

DEINE IDEE. DEIN IT-PROJEKT. DEINE ZUKUNFT.

Software Campus-Summit 2015

Software Campus-Summit 2015

Liebe TeilnehmerInnen und Gäste des zweiten Software Campus-Summits,

die IT-Branche leidet seit langem unter einem Fachkräftemangel. Es fehlen nicht nur IT-Experten, sondern vor allem auch Führungskräfte. Das Bundesministerium für Bildung und Forschung hat gemeinsam mit der Industrie den „Software Campus“ zur Förderung des Führungskräftenachwuchses konzipiert und gestartet. Unternehmen und Bundesregierung investieren gemeinsam und zu gleichen Teilen in die Zukunft unseres Landes.

Im Projekt Software Campus werden pro Jahr etwa 100 fortgeschrittene Informatikstudierende ausgebildet. Ihr Lernort ist die Praxis. Die Studierenden bearbeiten selbst gewählte Projekte unter persönlicher Anleitung eines Spitzenforschers und eines erfahrenen Managers aus IT-Unternehmen. Mit dieser Qualifizierung geben wir Studierenden das Rüstzeug auf ihrem Weg zu Spitzenpositionen in Wirtschaft und Forschung.

Der Software Campus verleiht in diesem Jahr neuen Absolventinnen und Absolventen das Abschlusszertifikat. Dabei zeigt sich, welche fachliche und persönliche Entwicklung die Studierenden genommen haben und für welche Positionen in der Wirtschaft einige der Absolventen bereits vorgesehen sind. Dieser Bedarf aus der Praxis ist der beste Beweis für die hohe Qualität des Software Campus.

Prof. Dr. Wolf-Dieter Lukas

(Ministerialdirektor im Bundesministerium für Bildung und Forschung)

-
- S. 4 **AdAPT:** Ein universeller Ansatz zur inkrementellen Aktivitätserkennung und -vorhersage in Smart Microgrids (*Jochen Frey*)
- S. 5 **AFTOS:** Spezifikation und Analyse der Verfügbarkeit Fehler-toleranter Systeme (*Maximilian Junker*)
- S. 6 **AKTIFUNK:** Untersuchung deterministischer Funkausbreitungsmodelle zur gerätfreien, funkbasierten Aktivitätserkennung (*Markus Scholz*)
- S. 7 **ANSE:** Analyse von Nutzungsdaten in der Software Evolution (*Sebastian Eder*)
- S. 8 **AUNUMAP:** Automatische Nutzercharakterisierung für Marktforschung und Prototypenentwicklung anhand psychographischer Daten aus Social Media und Sprachanwendungen (*Tim Polzehl*)
- S. 9 **BoRoVo:** Sichere, robuste, effiziente und flexible, dezentrale Wahlsysteme für spontane Wahlen in Managements (*Stephan Neumann*)
- S. 10 **CAMSA:** Konfigurierbare und überwachbare Dienstqualität für mandantenfähige Software-as-a-Service-Anwendungen (*Tilman Kopp*)
- S. 11 **CME:** Concurrent Manufacturing Engineering (*Anthony Anjorin*)
- S. 12 **CrowdMAQA:** Cloud Computing Location Metadata - Persönlicher Datenschutz und eine vertrauensvolle Privatsphäre durch selbstbestimmtes Cloud-Computing (*Sebastian Zickau*)
- S. 13 **CurCUMA:** Motivation and Automatic Quality Assessment in Paid Crowdsourcing Online Labor Markets (*Babak Naderi*)
- S. 14 **DaaS+:** Display as a Service Plus (*Alexander Löffler*)
- S. 15 **DuRMaD:** Dual Reality Management Dashboard für den Einzelhandel (*Gerrit Kahl*)
- S. 16 **EEinIR:** Balancing exploration and exploitation (multi armed bandit) in information retrieval systems (*Weijia Shao*)
- S. 17 **EHAC:** Erkennung Hardware-basierter Angriffe auf Computerhauptspeicher (*Patrick Stewin*)
- S. 18 **EnerFlow:** Verfahren für ein intelligentes Energiemanagement zur Optimierung des Energieverbrauchs in intelligenten Umgebungen unter Einbezug von Kontextinformationen (*Frank Englert*)
- S. 19 **Gaming-QoE-Model:** Mobiler-Campus-Charlottenburg (App): Plattform zur Untersuchung von QoE of Mobile Gaming (*Justus Beyer*)
- S. 20 **HDBC:** Hauptspeicherdatenbanken in der Cloud (*Tobias Mühlbauer*)
- S. 21 **IndRe:** Organisation von Software-Entwicklungsartefakten zur Verbesserung von Wiederverwendung im Kontext der industriellen Software-Entwicklung (*Veronika Bauer*)
- S. 22 **Intellektix:** Intelligente Informationsextraktion für das Entdecken und Verknüpfen von Wissen (*Sebastian Krause*)
- S. 23 **ISSA:** Inferenzbasierte Suche auf Supportanfragen (*Kathrin Eichler*)
- S. 24 **IT-GiKo:** Gitterbasierte Kryptographie für die Zukunft (*Rachid El Bansarkhani*)

-
- S. 25 **IUPD:** Immediate Usability für Interaktive Public Displays (*Robert Walter*)
- S. 26 **KaVE:** Verfeinerte Extraktion von Best-Practices aus Software Repositories durch Kombination von automatisierten Verfahren mit Expertenwissen (*Sebastian Proksch*)
- S. 27 **KoBeDie:** Kosteneffiziente Bereitstellung Cloud-basierter Unternehmensanwendungen unter Berücksichtigung von Dienstgüteanforderungen (*Ronny Hans*)
- S. 28 **KoSIUX:** Kontextsensitivität zur Steigerung der Sicherheit und User Experience mobiler Anwendungen (*Christian Jung*)
- S. 29 **LD-Cubes:** Linked Data Cubes – Integration und Analyse Verteilter Statistischer Datensätze durch Linked Data (*Benedikt Kämpgen*)
- S. 30 **ModSched:** Domänen-Unabhängiges Modulares Scheduler-Framework für Heterogene Multi- und Many-Core Architekturen (*Anselm Busse*)
- S. 31 **Move4Dynamic:** Modulares Vertrauensmanagement für dynamische Dienste (*Sascha Hauke*)
- S. 32 **OptimusPlant:** Dezentrale Steuerung von cyber-physischen Produktionssystemen (*Julius Pfrommer*)
- S. 33 **PEAKS:** Plattform zur effizienten Analyse und sicheren Komposition von Softwarekomponenten (*Ben Hermann*)
- S. 34 **Poker:** Policy-getriebene Konextualisierung in ereignisbasierter Middleware (*Tobias Freudenreich*)
- S. 35 **PreTIGA:** Optimierung von gamifizierten Anwendungen durch Echtzeitvorhersagen von Nutzerverhalten und Anreizen (*Stefan Tomov*)
- S. 36 **RefMod@RunTime:** Induktive Entwicklung von Referenzmodellen mithilfe von Maschinensemantik und deren praktische Evaluation (*Jana Rehse*)
- S. 37 **ReproFit:** Automatische Extraktion von Reproduktionsschritten für Software-Fehler aus Benutzer-Interaktionen (*Tobias Röhm*)
- S. 38 **SatIN:** Verarbeitung von inkongruenten Intentionen in systemgestützten, non-kollaborativen Dialogen (*Sabine Janzen*)
- S. 39 **SCORE:** Collaborative Enrichment of Business Processes (*Christina Di Valentin*)
- S. 40 **SiPri:** Situationsbedingte Privacy Interfaces (*Florian Gall*)
- S. 41 **SMeC:** Secure Media Cloud (*Alexandra Mikityuk*)
- S. 42 **SUITE:** Semi-automatische Erkennung und Wissensbasis-Integration von Neuigkeiten aus Textdokumenten (*Michael Färber*)
- S. 43 **USV:** Usability von Software-Verifikationssystemen: Evaluierung und Verbesserung (*Sarah Grebing*)
- S. 44 **WEFITM:** Weiterentwicklung der FIT-Metrik für Ereignis-basierte Softwaresysteme (*Sebastian Frischbier*)

Ein universeller Ansatz zur inkrementellen Aktivitätserkennung und -vorhersage in Smart Microgrids

Jochen Frey

AdAPT

Das Verstehen menschlicher Handlungsweisen und der dahinter liegenden Intention nimmt eine immer wichtigere Rolle bei der Entwicklung von intelligenten Benutzerschnittstellen in unterschiedlichen Anwendungsfeldern ein. Zum Beispiel hat der Bereich Ambient Assisted Living (AAL) das Ziel, vor allem ältere Menschen bei der Ausführung von Aktivitäten des täglichen Lebens (ATL) zu unterstützen und gegebenenfalls externe Hilfe anzufordern. Durch die Auswertung von Sensordaten können Smart Homes automatisch erkennen, welche Aktivitäten von ihren Bewohnern aktuell ausgeführt werden und ob es zu Problemen oder Abweichungen innerhalb ihrer normalen Tagesroutine kommt und daraufhin unterstützend wirken.

Im Bereich sogenannter Smart Microgrids finden Methoden zur Aktivitätserkennung eine sinnvolle Anwendung. So können zum Beispiel zukünftige Stromverbrauchsdaten der im intelligenten Stromnetz angeschlossenen Smart Homes auf Basis von erlernten und individuellen Aktivitätsprofilen der Bewohner besser prognostiziert werden. Die Energieerzeuger werden hierdurch in die Lage versetzt, detaillierte Lastprofile zu erstellen und somit ihre Stromkontingente am Energiemarkt effizienter und kostengünstiger zu handeln.

Das übergeordnete Ziel des Projektes AdAPT ist die Realisierung eines universellen Ansatzes zur Erkennung von Aktivitäten innerhalb eines Smart Homes. Die Grundidee ist hierbei, dass sich die intelligente Umgebung dynamisch und kontinuierlich auf die Benutzer und die sich ändernden Gegebenheiten anpasst. Hierzu wird im Projekt ein verteiltes Plattformkonzept entwickelt und realisiert, welche die Integration von neuartigen digitalen Mehrwertdiensten (Smart Services) in eine intelligente Umgebung ermöglicht. Darüber hinaus werden sowohl ein Marktplatzportal zur Verbreitung von Smart Services als auch eine durchgängige Werkzeugkette zur Entwicklung von Smart Services bereitgestellt. Abschließend werden, aufbauend auf dem vorgestellten Konzept, zwei unterschiedliche Smart Services zur Erkennung von Benutzeraktivitäten innerhalb eines Smart Homes implementiert.

Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und Deutsche Telekom AG

Projektlaufzeit: 01.03.2013 bis 28.02.2015

Spezifikation und Analyse der Verfügbarkeit Fehler-toleranter Systeme

Maximilian Junker

AFTOS

Das Paradigma Modellbasierung trägt heute schon in vielen Bereichen der Software- und Systementwicklung dazu bei, die Effizienz der Entwicklung und die Qualität der Ergebnisse zu steigern. Viele wichtige Eigenschaften des Systems können schon in frühen Phasen überprüft werden und ihre weitere Einhaltung in späteren Phasen sichergestellt werden.

Ziel des Software Campus-Projekt AFTOS war es, die Prinzipien der modellbasierten Entwicklung von rein funktionalen Eigenschaften auf Verfügbarkeitseigenschaften auszuweiten. Gerade bei eingebetteten Systemen ist eine hohe Verfügbarkeit von großer Bedeutung. Beispiele für Systeme, bei denen Verfügbarkeit eine besonders große Rolle spielt, sind Steuerungssysteme für Versorgungsnetze (Energie und Telekommunikation) und Transportsysteme.

Die konkreten Inhalte des Forschungsprojekts sind

1. Eine Interviewstudie, bei der wir insgesamt 15 Verfügbarkeitsexperten aus verschiedenen Unternehmen der deutschen Industrie befragt haben. Ergebnisse dieser Studie waren u.a. eine Darstellung des Standes der Technik bzgl. Verfügbarkeit heute sowie die Identifikation von konkreten Problemen in Bezug auf eine systematische Erreichung von Verfügbarkeit.
2. Eine Modellierungstechnik und -methodik, mit der sich Verfügbarkeitsanforderungen und relevante Systemeigenschaften modellieren lassen. Auf dieser Basis können Verfügbarkeitsanforderungen formuliert und ihre Erfüllung analysiert werden.
3. Werkzeugunterstützung für die Modellierung und Analyse, basierend auf einem Open Source Modellierungswerkzeug.
4. Eine Fallstudie, bei der die entwickelten Techniken auf den Teil eines Zugsteuerungssystem angewandt wurden.

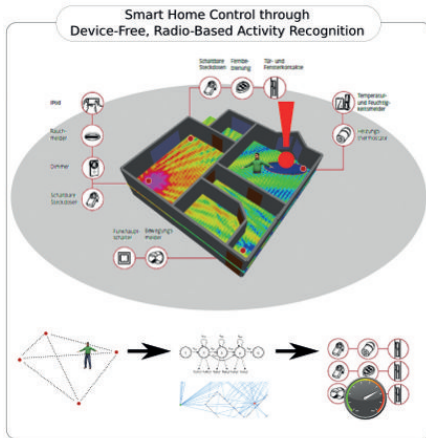
Projektpartner: Technische Universität München und Siemens AG

Projektlaufzeit: 01.01.2013 bis 31.03.2015

Untersuchung deterministischer Funkausbreitungsmodelle zur gerätefreien, funkbasierten Aktivitätserkennung

Markus Scholz

AKTIFUNK



Dem Markt für die Hausautomatisierung wird weltweit ein starkes Wachstumspotenzial prognostiziert. Eine der Herausforderungen der Hausautomatisierung ist die Kontrolle der vielen vernetzten, intelligenten Geräte, die in der Wohnung der Zukunft zu finden sind. Aktuelle Produkte setzen auf Eingaben mittels klassischer Benutzerschnittstellen (Tablet, Smartphone, etc.). Mit steigender Anzahl von Geräten benötigt diese Art der Bedienung mehr Zeit und reduziert daher die Nutzerakzeptanz. Eine implizite Steuerung, bei der im Hintergrund die Aktivitäten des Benutzers erfasst und analysiert und diese in Einstellungen der Hauselektronik umgesetzt werden, könnte dieses Problem lösen.

Ein implizites Automatisierungssystem kann aus zwei Komponenten bestehen: Einer Sensorikkomponente zur Erfassung der aktuellen Benutzersituation und einer Steuerungskomponente welche die Sensorikinformationen für die Steuerung der Hausgeräte einsetzt. In AKTIFUNK wird eine neue Sensorikkomponente entwickelt die eine akzeptable und kosten-effektive Erfassung von Benutzeraktivitäten ermöglicht. Dabei werden Signalqualitätsinformation von Funksystemen für die Erkennung von menschlichen Aktivitäten eingesetzt. Im Projekt wird ein Prototyp dieser Sensorikkomponente in Hard- und Software entwickelt. Da Radiowellen nicht vom Menschen wahrgenommen werden und viele Smart-Home-fähige Geräte mit einer Funkschnittstelle ausgestattet sind, wird dadurch ein hoher Komfort bei der Interaktion gewährleistet und die Installation einer zusätzlichen Infrastruktur kann vermieden werden.

Projektpartner: Karlsruher Institut für Technologie (KIT) und Robert Bosch GmbH

Projektlaufzeit: 01.01.2014 bis 30.04.2015

Analyse von Nutzungsdaten in der Software Evolution

Sebastian Eder

ANSE



Softwaresysteme werden entwickelt, um einen bestimmten Zweck zu erfüllen. Erst wenn ein Softwaresystem ausgeführt wird, erfüllt es seinen Zweck und entwickelt seinen Nutzen. Dadurch entsteht der Geschäftswert des Softwaresystems. Allerdings können sich der Zweck des Softwaresystems oder die Bedürfnisse von Benutzern mit der Zeit ändern. Dies ist besonders bei langlebigen Softwaresystemen der Fall.

Um eine Änderung im Zweck eines Softwaresystems oder in den Bedürfnissen der Benutzer zu erkennen, kann eine Nutzungsdatenanalyse hinzugezogen werden, bei der die tatsächliche Benutzung eines Softwaresystems gemessen wird. Dadurch wird explizit, welche Teile eines Softwaresystems tatsächlich genutzt werden, und welche nicht. Nutzung ist ein Indikator für wichtige Systemteile und fehlende Nutzung zeigt an, dass ein Systemteil möglicherweise unwichtig ist.

Da die Wartung einen großen Teil der Gesamtkosten eines langlebigen Softwaresystems verschlingt, ist besonders hier Optimierungspotenzial gegeben. Mit der Analyse von Nutzungsdaten, und somit mit der Einstufung von Programmteilen auf Relevanz, können Wartungsarbeiten besser priorisiert und koordiniert werden, um Kosten im Wartungsprozess einzusparen.

In ANSE werden Techniken zur Analyse der tatsächlichen Nutzung von Softwaresystemen entwickelt. Dabei werden nicht nur Daten der Nutzer eines Softwaresystems analysiert, sondern auch Daten aus der Testumgebung eines Systems. Auf diese Weise können nicht nur Entwicklungsaufwände, sondern auch Testaufwände während der Softwareevolution besser priorisiert und koordiniert werden.

Projektpartner: Technische Universität München und Siemens AG

Projektlaufzeit: 01.02.2013 bis 31.03.2015

Automatische Nutzercharakterisierung für Marktforschung und Prototypenentwicklung anhand psychographischer Daten aus Social Media und Sprachanwendungen

Tim Polzehl

AUNUMAP

Automatische Nutzercharakterisierung für Marktforschung und Prototypenentwicklung anhand psychographischer Daten aus Social Media und Sprachanwendungen wie Facebook, Apple, Google und Twitter machen es uns vor: Neuartige Nutzeranalysen und Analysen von Nutzungsverhalten bergen ein enormes Potenzial für einzigartige Rückschlüsse auf die Bedürfnisse der Nutzer. Im Mittelpunkt der Nutzeranalysen stehen die sogenannten psychographischen Daten, die sehr viele Nutzer im täglichen Umgang mit Sozialen Netzwerken oder durch die immer populärer werdende Möglichkeit der Sprachinteraktion hinterlassen. Meinungen, Vorschläge, Anliegen und Beschwerden werden millionenfach täglich gepostet – eine Datenmenge, die ohne automatische Verfahren nicht mehr analysierbar ist.

Ziel des Projektes AUNUMAP ist es, passgerecht und tagesaktuell die Bedürfnisse einer großen Anzahl von Nutzern automatisch zu erkennen und diese Informationen möglichst nahtlos Marketing- und/oder Produktentwicklungsprozessen zugänglich zu machen. Die psychographischen Eigenschaften sind Ziele der automatischen Klassifikation im vorgeschlagenen Projekt. Eine Software oder ein Service zur Datenaufbereitung hinsichtlich dieser Eigenschaften ist derzeit nicht existent. Eine Extraktion und Aufbereitung dieser Nutzerinformationen stellt somit eine Innovation dar. Diese neuartigen Nutzeranalysen bürden ein enormes Potenzial für einzigartige Rückschlüsse auf eine Charakterisierung und Bedürfnisanalyse der Nutzer.

Am Ende des Projektes sollen Algorithmen und Prototypen stehen, die für Analysen von Social-Media-Daten und Sprachdaten eingesetzt werden können. Des Weiteren werden diese Algorithmen und Resultate wichtige Parameter und Nutzereigenschaften aufzeigen, die im gesamten wissenschaftlichen Umfeld der Mensch-Maschine-Interaktion von fundamentalem Interesse sind. Mit den beschriebenen Eigenschaften lassen sich Schnittstellen gestalten, die zu einem intuitiveren Benutzererlebnis und zu einer natürlicheren Mensch-Maschine-Kommunikation beitragen, bspw. beim Einsatz in Dialogsystemen, die somit Stimmungen und persönliche Wertevorstellungen der Nutzer erkennen, und darauf Rücksicht nehmen können. Ergebnisse und Algorithmen werden ebenfalls helfen vorhandene Schnittstellen dynamischer und adaptiver zu gestalten.

So könnten sich bspw. Maschinen an den Erfahrungsgrad ihrer Nutzer anpassen, bspw. bei Eingabe einer Websuche oder beim Aufbereiten der Ergebnisse einer solchen.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.03.2013 bis 28.02.2015

Sichere, robuste, effiziente und flexible, dezentrale Wahlsysteme für spontane Wahlen in Managements

Stephan Neumann

BoRoVo

Seit einiger Zeit steigt das Interesse an Internetwahlsystemen. So werden Internetwahlsysteme nicht mehr nur für Wahlen in Vereinen, Universitäten und Betrieben eingesetzt, sondern auch für parlamentarische Wahlen und Abstimmungen wie zum Beispiel in Estland und in der Schweiz. Basierend auf einer breiten Palette kryptographischer Primitive und Protokolle garantieren Internetwahlsysteme eine Reihe von Sicherheitseigenschaften, wovon die Wichtigsten die Geheimhaltung der Wahl sowie die öffentliche Nachvollziehbarkeit der korrekten Ergebnisermittlung sind. Hierzu verteilen die eingesetzten Internetwahlsysteme Zuständigkeiten und Vertrauen auf verschiedene Wahlautoritäten wie Wahlserver, Administratoren und Schlüsselinhaber. Durch den damit verbundenen finanziellen und administrativen Aufwand sind diese Internetwahlsysteme für Wahlen, in denen sich der Bedarf einer Entscheidungsfindung spontan und verteilt ergibt, nicht unmittelbar geeignet. Entsprechend fehlen geeignete Internetwahlsysteme für beispielsweise Wahlen und Abstimmungen in Gremien und Aktiengesellschaften.

Im Rahmen des durchgeführten Vorhabens wurde ein entsprechendes Internetwahlsystem als dezentrales Wahlsystems für mobile Endgeräte entwickelt. Hierzu wurde in einem ersten Schritt eine Anforderungsanalyse durchgeführt. Dabei wurden rechtliche Anforderungen in technische Anforderungen überführt sowie Anforderungen von Endanwendern identifiziert. In einem zweiten Schritt wurde die Konzeptentwicklung des dezentralen Wahlsystems auf Grundlage existierender Literatur durchgeführt. Konzeptuell basiert das Wahlsystem auf einer verteilten kryptographischen Schlüsselerzeugung und Entschlüsselung, sowie einer verifizierbaren Mix-basierten Stimmen-anonymisierung. Nach dem erfolgreichen Abschluss der Konzeptentwicklung sowie dessen Evaluation gegen die Anforderungen wurde die Implementierung des Konzepts als Smartphone Applikation umgesetzt. Die Implementierung wurde mithilfe etablierter kryptographischer Bibliotheken durchgeführt. Im letzten Teil des Projektes wurde die Benutzbarkeit der entwickelten Applikation in einer Benutzerstudie evaluiert. Die Ergebnisse dieser Studie belegen die Benutzbarkeit der entwickelten Applikation.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.02.2013 bis 30.11.2014

Konfigurierbare und überwachbare Dienstqualität für mandantenfähige Software-as-a-Service-Anwendungen

Tilman Kopp

CAMSA

Um im globalen Wettbewerb konkurrenzfähig bleiben zu können, konzentrieren sich erfolgreiche Unternehmen auf die Ausführung ihrer wertschöpfenden Prozesse. Damit diese effizient betrieben werden können, setzen Unternehmen Software zur Prozessausführung und Ressourcenverwaltung ein. Der eigenständige Betrieb und die Instandhaltung der Unternehmenssoftware tragen jedoch nicht unmittelbar zum Unternehmenserfolg bei. Um Kosten zu sparen lagern daher zunehmend klein- und mittelständische Unternehmen den Betrieb ihrer Unternehmenssoftware aus. Eine zunehmende Form der Auslagerung stellt die Verwendung von Software-as-a-Service (SaaS) Lösungen dar. Sie ermöglicht die Nutzung von Geschäftsanwendungen über das Internet. Da SaaS-Anbieter viele Kunden mittels derselben technischen Infrastruktur bedienen und somit ihre Ressourcen besser auslasten, können Kostenvorteile realisiert und an die Kunden weitergegeben werden.

Mit der gemeinsamen Nutzung der technischen Infrastruktur gehen jedoch auch neue Probleme einher. So unterscheiden sich die Anforderungen von Kunden gerade bei komplexen Unternehmenslösungen, wie beispielsweise im Bereich der Unternehmensressourcenplanung. In Bezug auf funktional unterschiedliche Anforderungen einzelner Kunden bieten moderne SaaS-Anwendungen daher kundenspezifisch konfigurierbare Funktionalität. Eine kundenspezifische Differenzierung von Qualität und deren kunden- und anbieterseitige Überwachung ist in vielen SaaS-Produkten hingegen bislang kaum ausgeprägt. Im Rahmen des Projekts „Konfigurierbare und überwachbare Dienstqualität für mandantenfähige Software-as-a-Service Anwendungen“ (CAMSA) wurde daher untersucht, wie sich kundenspezifische Qualitätsanforderungen hinsichtlich der Anwendungsperformanz durch geeignete Softwarearchitekturen gewährleisten lassen. Anhand eines definierten Szenario- und Prozessmodells wurden zunächst existierende Architekturen und zur Beeinflussung der Performanz geeignete Verfahren analysiert. Dabei haben sich Architekturen, welche eine dynamische Partitionierung der Anwendungen ermöglichen, als Lösung empfohlen. Im Rahmen von CAMSA wurden diese Ansätze mit dem Ziel der Reduktion des Synchronisationsaufwandes zwischen Komponenten einer Partition erfolgreich weiterentwickelt. Die optimierten Ansätze wurden schließlich auf ihre Anwendbarkeit hin untersucht und die Ergebnisse veröffentlicht.

Projektpartner: Karlsruher Institut für Technologie (KIT) und SAP SE

Projektlaufzeit: 01.03.2013 bis 28.02.2015

Concurrent Manufacturing Engineering

Anthony Anjorin

CME

In modernen Ingenieurverfahren ist es nötig, mit untereinander abhängigen Artefakten unterschiedlicher Art zu arbeiten. Ingenieure greifen dabei typischerweise gleichzeitig und mit verschiedenen Werkzeugen auf diese Artefakte zu. Dieses heterogene Vorgehen birgt die Gefahr von Inkonsistenzen zwischen den beteiligten Artefakten.

In der Fertigungstechnik, ein für ein Exportland wie Deutschland sehr wichtiger Industriezweig, ist es besonders wichtig, dieser Herausforderung der Konsistenzhaltung zu begegnen, da eine ganze Kette von Werkzeugen mit eigenen, zum Teil normierten (d.h. durch einen Standard vorgegeben und nicht mehr zu ändernden), Formaten verwendet werden.

Obwohl die Vorwärtsrichtung durch die Werkzeugkette meistens gut unterstützt ist (Stand der Technik), ist aus Erfahrungen in der Softwareentwicklung klar, dass eine iterative Vorgehensweise unabdingbar ist, um Kosten zu sparen und die Effizienz der Fertigungsprozesse zu steigern. Es muss daher möglich sein, in jedem Schritt des Prozesses Änderungen am aktuellen Artefakt auch rückwärts durch die Kette zu propagieren. Mit Ansätzen und Technologien aus der Metamodellierung ist es möglich, für die Konsistenzaufgabe relevante Aspekte der beteiligten Artefakte als formale Modelle zu spezifizieren, was eine Beschreibung der Transformationen mittels Modelltransformationssprachen ermöglicht. Insbesondere können bidirektionale Modelltransformationssprachen durch eine gezielte Propagation inkrementeller Änderungen, die Konsistenz zwischen Artefakten sicherstellen.

Im Rahmen des IT-Forschungsprojekts „Concurrent Manufacturing Engineering“ (CME) wurde diese Konsistenzaufgabe im Bereich der Fertigungstechnik am Beispiel einer industriellen Fallstudie untersucht. Das Hauptziel des Projekts war es, durch die Untersuchung einer neuen Anwendungsdomäne mit neuen speziellen Anforderungen, neue Erkenntnisse, Ideen und Denkanstöße für Sprachfeatures und Forschungsrichtungen im Bereich von bidirektionalen Modelltransformationssprachen zu liefern.

Ergebnisse waren unter anderem ein neues Modularitätskonzept, um die Komplexität größerer Transformationen zu beherrschen, sowie eine Reihe von statischen Analysemethoden, um Spezifikationsfehler frühzeitig und automatisiert aufzudecken. Durch die Zusammenarbeit mit einem Industriepartner konnte das Potenzial und die Praxistauglichkeit aktueller Forschungsideen und Prototypen im Bereich bidirektionaler Modelltransformationen anhand eines realistischen Anwendungsbeispiels mit industrieller Relevanz evaluiert werden.

Projektpartner: Technische Universität Darmstadt und Siemens AG

Projektlaufzeit: 01.01.2013 bis 31.12.2013

Motivation and Automatic Quality Assessment in Paid Crowdsourcing Online Labor Markets

Babak Naderi

CrowdMAQA



Crowdsourcing auf Basis von Microwork-Web-Plattformen eröffnet einmalige Vorteile, um große und spezifische Daten für die Marktforschung bzw. für Entwicklung und Forschung zu generieren. Hauptvorteile sind die hohe Geschwindigkeit der Erhebung, die flexible Ausgestaltung und die geringen Kosten. Gleichzeitig können die generierten Daten aber auch einen höheren Rauschanteil aufzeigen, d.h. unter Umständen eine geringe Reliabilität oder Validität aufweisen. Die Sicherstellung einer geeigneten Datenqualität wird vor allem dann relevant, wenn die Arbeiter der Crowd dafür bezahlt werden.

Ausgehend von der wachsenden Verfügbarkeit und steigenden Popularität von Crowdsourcing-Daten will das Projekt CrowdMAQA zentrale Fragen zu Qualität und Eigenschaften mittels Crowdsourcing erhobener Daten beantworten. Das Projekt konzentriert sich dabei auf die Beurteilung, Vorabschätzung, Überwachung und Sicherstellung der resultierenden Qualität und untersucht Faktoren, die die Motivation der Arbeitnehmer beeinflussen. Um solche Qualitätseinbußen auszuschließen werden im vorgeschlagenen Projekt automatische Qualitätssicherungsmodelle konzipiert, implementiert und evaluiert. Solche Modelle können bspw. die Reihenfolge der Aufgaben für die Befragten auswählen oder die Qualität der generierten Daten anhand von Konsistenten und/oder geloggtten Zeitangaben vorhersagen und somit über die Bezahlung mitentscheiden. Darüber hinaus sollen Faktoren ermittelt werden, die die Motivation und deren Einfluss auf die Arbeitsleistung ausdrücken. Dabei ist zu berücksichtigen, dass die Qualität der Arbeit und die Motivation des Arbeiters (Probanden) von einer ganzen Reihe verschiedener Faktoren beeinflusst werden können.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.04.2013 bis 31.03.2015

Cloud Computing Location Metadata – Persönlicher Datenschutz und eine vertrauensvolle Privatsphäre durch selbstbestimmtes Cloud-Computing

Sebastian Zickau

CurCUMA

Die Nutzung von leistungsfähigen mobilen Endgeräten und Cloud Computing-Diensten ist für viele Menschen heute ein Teil des Alltags. Die Daten-Typen, die bei der Nutzung anfallen und verarbeitet werden, sind in den vergangenen Jahren vielseitiger geworden. Es werden u.a. Dokumente, Bilder und Videos über Internetdienste verschickt bzw. mit anderen Nutzern geteilt. Die Natur digitaler Daten führt auch zu Meta-Informationen, die im Lebenszyklus der in der Cloud gespeicherten Daten anfallen. Mit Meta-Informationen bezeichnet man allgemein Daten über Daten.

Das Curcuma-Projekt nutzt spezifische, dynamische und semantische Meta-Informationen und Methoden, um dem Nutzer verschiedene Mehrwertdienste anzubieten. Sicherheitsfunktionalitäten stehen, gerade vor dem Hintergrund aktueller Pressemeldungen, auf der Wunschliste vieler Anwender. Eine mobile Anwendung erlaubt es dem Besitzer digitaler Daten diese, mit dynamischen und nicht-dynamischen Attributen (Meta-Informationen) zu versehen, um z.B. die Zugriffe nur von einem bestimmten Ort aus zu erlauben. Bei besonders sensiblen Informationen, z.B. medizinische Daten, wird als Information festgehalten, wer wann darauf zugegriffen hat, d.h. der Zeitpunkt und der Ort des Zugriffs des behandelnden Arztes.

Ein Fokus des Projekts liegt auf der Nutzung von dynamischen und nicht-dynamischen Ortsinformationen. Dynamisch sind diese Informationen, wenn der Zugriff von mobilen Endgeräten erfolgt. Nicht dynamisch sind die Ortsinformationen, wenn Server-Standorte von Internet- und Cloud Computing-Diensten mit in den Prozess einfließen.

Für die nicht-dynamischen Attribute wird im Laufe des Projekts ein Webdienst eingerichtet und erweitert, mit dem sich diese Informationen auch visuell darstellen lassen. Moderne semantische Methoden helfen dabei, diese Informationen mit bereits bestehenden zu verknüpfen. Dies können u.a. Informationen zu rechtlichen Rahmenbedingungen und Datenschutzgesetzen in Bezug zu den Internet- und Cloud-Diensten sein. Der Web-Dienst des Unterprojekts trägt den Namen OpenCloudComputingMap. Allgemein gesagt, befasst sich das Vorhaben Curcuma mit der Erforschung und Entwicklung einer Lösung für einen selbstbestimmten und sicheren Umgang mit persönlichen Daten im Internet.

Projekt-Webseite:
<http://www.curcuma-project.net/>

Unterprojekt-Webseite:
<http://www.opencloudcomputingmap.org/>

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.03.2014 bis 29.02.2016

Display as a Service Plus

Alexander Löffler

DaaS+



Display as a Service (DaaS) ist eine Technologie, die auf die Virtualisierung der Verbindung zwischen Pixelquellen und Displays abzielt. Anstelle üblicher, dedizierter 1-zu-1-Verbindungen über Videokabel (z.B. HDMI) sind Quellen und Bildschirme nur noch Ressourcen im überall verfügbaren, drahtgebundenen oder drahtlosen Netzwerk (Internet), und kommunizieren flexibel in einer m-zu-n-Beziehung. Im Software Campus-Projekt „DaaS Plus (DaaS+)“ wurde die Technologie nun einerseits technisch weiterentwickelt und andererseits echten Benutzern zur Verfügung gestellt, deren Erfahrungen mit DaaS+ wiederum in die Weiterentwicklung zurückflossen.

So wurde in DaaS+ bspw. das quellseitige Video-Encoding in Echtzeit durch die Verwendung der GPU-Hardware in aktuellen Intel-Prozessoren um Faktor 2-3 beschleunigt, während gleichzeitig die CPU entlastet wird. Dies ermöglicht wiederum die Verwendung noch kleinerer und günstigerer Geräte an allen Enden des Systems. Zur Evaluierung durch Endnutzer wurden zwei Use-Cases beim T-Systems Innovation Center München (ICM) als "prototypischer Industriekunde" umgesetzt. Die Szenarien "Kollaboratives Meeting" und "Virtuelles Fenster ins Rechenzentrum" fokussierten sich auf die Flexibilität der Pixelverteilung von mehreren mobilen Quellen auf mehrere Displays bzw. auf die dynamisch gesteuerte, synchrone Ausgabe auf eine Display-Wall.

Die DaaS-Technologie wurde während der Laufzeit von DaaS+ mit dem ersten Preis des CeBIT Innovation Award (2013) ausgezeichnet und seitdem in einer Vielzahl industrieller Projekte eingesetzt. Anwendungen reichen dabei von einfachen Präsentationen im Besprechungsraum über In-Car-Szenarien bis hin zu High-End-Visualisierungen in industriellen Visualisierungszentren.

Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und Deutsche Telekom AG

Projektlaufzeit: 01.03.2013 bis 28.02.2015

Dual Reality Management Dashboard für den Einzelhandel

Gerrit Kahl

DuRMaD

Im Projekt DuRMaD wurde eine Möglichkeit der visuellen Darstellung des aktuellen Zustands von instrumentierten Einkaufsumgebungen und deren Sensorinformationen mittels eines interaktiven 3D-Modells realisiert. Mit dem entwickelten Dual Reality Management Dashboard (DuRMaD) wurde eine entsprechende Infrastrukturumgebung entwickelt, die es ermöglicht, alle relevanten Sensoren und Aktuatoren aus einer einzigen benutzerorientierten Bedienoberfläche heraus zu kontrollieren, zu analysieren und zu steuern. Die erfassten Daten werden hierbei in einem 3D-Modell der Umgebung entsprechend aufbereitet dargestellt.

Für die Aufbereitung der Daten wurde ein Verfahren entwickelt und implementiert, welches eine Erstellung von Business Intelligence (BI) Diensten ermöglicht. Diese BI-Dienste können verwendet werden, um die erfassten Informationen aufzubereiten und die Ergebnisse im 3D-Modell darzustellen. Des Weiteren wurde eine Kommunikationsinfrastruktur zum bi-direktionalen Informationsaustausch und der gegenseitigen Beeinflussung zwischen den beiden Welten geschaffen, was als Dual Reality (DR) bezeichnet wird.

Das in diesem Projekt entwickelte Framework unterstützt hierbei eine Visualisierung über eine Internetseite, indem eine Schnittstelle zu XML3D integriert wurde. XML3D wird an der Universität des Saarlandes entwickelt und hat den Anspruch, interaktive 3D-Inhalte ohne spezielle Plug-Ins in Webseiten darstellen zu können. DuRMaD kann durch einfache Schnittstellen, welche direkt in die Kommunikationsinfrastruktur zwischen den Systemen integriert sind, die Sensoren und Aktuatoren der Umgebung einbinden. Die Ergebnisse der Arbeit wurden in das Innovative Retail Laboratory (www.innovative-retail.de) integriert und sind dort ein wichtiger Bestandteil der vernetzten Umgebung geworden.

Zudem wurde das entwickelte System in die Demonstratortour durch das Labor integriert. Des Weiteren wurde eine automatische Floorplanning-Komponente implementiert, welche in Bezug auf positionsabhängige Verkaufsprognosen eine Regalbestückung mit möglichst optimalem Gewinn generiert. Hierzu müssen ausschließlich Produkte und Regalböden manuell selektiert werden. Darauf aufbauend werden mittels eines heuristischen Algorithmus verschiedene Regalbestückungen analysiert und bewertet. Das Ergebnis wird anschließend als Planogramm ausgegeben und kann Mitarbeitern als Vorräumungsvorgabe zur Verfügung gestellt werden.

Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und SAP SE

Projektlaufzeit: 01.01.2013 bis 30.06.2014

Balancing exploration and exploitation (multi armed bandit) in information retrieval systems

Weijia Shao

EEinIR

Viele Internet-Dienstleister müssen große Mengen von Objekten verwalten. Zum Beispiel Nachrichtenartikel, Werbungen, oder Geräte, die mit dem Internet-der-Dinge verbunden sind. Diese Objekte verändern sich häufig und dynamisch. Dadurch ist es schwierig, die Präferenzen der Nutzer ausschließlich mit Informationen zum Inhalt zu modellieren. Allerdings basieren die Geschäftsmodelle der Dienstleister stark auf der korrekten Vorhersage der Präferenzen. In der Praxis sammeln die Internet-Dienstleister deswegen zusätzlich das Feedback von ihren Kunden um die Präferenzen von den Kunden zu lernen (Exploration). Diese Erkenntnisse verwerten (Exploitation) sie, um Einkommen zu maximieren, woraus folgendes Problem entsteht:

- Falls die Internet-Dienstleister ihre Aufmerksamkeit auf das Sammeln des Feedbacks richten, kann ein kurzfristiger Einkommensverlust auftreten, wegen der hohen Wahrscheinlichkeit, den Kunden viele unerwartete Objekte anzubieten.
- Ein langfristiger Einkommensverlust kann auftreten, wenn die Internet-Dienstleister ihre Aufmerksamkeit auf das Maximieren des Einkommens (ohne weiteres Lernen der Präferenzen der Benutzer) richten und so den Kunden immer suboptimale Objekte anbieten.

Das Exploration-Exploitation-Problem (Multi-Armed-Bandit-Problem) wurde im Jahr 1930 das erste mal beschrieben, und spielt seitdem eine immer wichtigere Rolle für moderne Retrieval- und Empfehlungssysteme. Die vorhandenen Algorithmen haben drei Nachteile:

- Die vorhandenen Algorithmen basieren auf unrealistischen Annahmen.
- Die Komplexität der vorhandenen Algorithmen hängt von der Anzahl der Objekte ab. Deshalb sind sie ineffektiv für Systeme mit einer großen Menge von Objekten.
- Die vorhandenen Algorithmen sind schwer zu verwenden, da die manuellen Einstellungen des Eingabeparameters benötigt werden.

Außerdem sind die vorhandenen Computerprogramme, die Bandit-Algorithmen anbieten, domain-spezifisch, obwohl das Exploration-Exploitation-Problem domainunabhängig ist.

In diesem Projekt werden Algorithmen für das Exploration-Exploitation-Problem entwickelt, analysiert, und in ein von Data-Analytik-Anwendungen aufrufbares Framework eingebaut. Dieses Framework können Firmen aus allen industriellen Sektoren nutzen, die vor dem Multi-Arm-Bandit-Problem stehen, und damit durch die Algorithmen des Frameworks unterstützt werden.

Projektpartner: Technische Universität Berlin und SAP SE

Projektlaufzeit: Start in 2014

Erkennung Hardware-basierter Angriffe auf Computer Hauptspeicher

Patrick Stewin

EHAC

Das EHAC-Projekt nimmt sich den neuesten IT-Sicherheits-Herausforderungen, die durch heimliche Angriffe auf Computersysteme hervorgerufen werden, an. Moderne bösartige Software versteckt sich in Peripheriegeräten (z.B. Netzwerkkarte), um unautorisiert auf den Hauptspeicher der Computerplattform zuzugreifen. Im Hauptspeicher befinden sich zur Laufzeit viele sensible Daten, wie kryptografische Schlüssel oder geheime Dokumente, die vor Angriffen geschützt sein sollten. Ein Angriff erfolgt über Speicherdirektzugriff, d.h. der Hauptprozessor ist unbeteiligt. Bösartige Speicherdirektzugriffe waren vor dem EHAC-Projekt sogar unsichtbar für den Hauptprozessor.

Antivirensoftware kann solche hardwarebasierten Trojaner bzw. Rootkits nicht finden, da sie in der Regel die separierten Ausführungsumgebungen der Peripheriegeräte nicht berücksichtigen. Präventive Maßnahmen, wie hardwarebasierte Speicherzugriffskontrolle, gelten nicht notwendigerweise als ausreichend. Um auf den Computerhauptspeicher zuzugreifen, nutzt der Schädling den Systembus. Solche Zugriffe können zwar durch zusätzliche Hardwarefeatures eingeschränkt werden. In der Praxis werden solche Features jedoch nicht ausreichend unterstützt. Somit eignen sich diese Schädlinge ideal für Advanced Persistent Threats.

Zur Abwehr wurde der Machbarkeitsnachweis für einen neuartigen Detektor zur Erkennung von Hardware-basierten Angriffen auf Computerhauptspeicher – EHAC – entwickelt. Der Machbarkeitsnachweis basiert auf folgender Idee. Das Betriebssystem weiß über sämtliche Ein- und Ausgaben (Festplattenzugriffe, Tastatureingaben) der Peripherie Bescheid. Wenn darüber hinaus Daten transferiert werden, kann es sich nur um einen Angriff handeln. Diese Idee lässt sich zu folgender EHAC-Forschungsfrage konkretisieren: Lassen sich vom Hauptprozessor Systembusdatenpakete erkennen, die nicht zur Verarbeitung vom Hauptprozessor gedacht sind (eingehende Daten) oder gedacht waren (ausgehende Daten)?

Zur Erkennung unzulässig transferierter Daten nutzt der EHAC-Detektor den Fakt, dass sich der Hauptprozessor und die Peripheriegeräte den Systembus teilen. Zugriffe auf diese Ressource können jedoch nicht absolut zeitgleich erfolgen, sondern müssen arbitriert werden. Dadurch ergeben sich Seiteneffekte, die sich mit entsprechenden Messmechanismen erfassen lassen. Der EHAC-Detektor implementiert entsprechende Messmechanismen zur Laufzeitüberwachung des Systembusses und kann somit die oben beschriebenen Angriffe zuverlässig erkennen.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.01.2013 bis 30.06.2014

Verfahren für ein intelligentes Energiemanagement zur Optimierung des Energieverbrauchs in intelligenten Umgebungen unter Einbezug von Kontextinformationen

Frank Englert

EnerFlow

Energiesparmaßnahmen und energieeffizientes Verhalten gewinnen aufgrund steigender Energiepreise und knapper werdenden Ressourcen zunehmend an Bedeutung. Immer mehr Privathaushalte und Unternehmen versuchen daher, ihren Energieverbrauch zu reduzieren. Oftmals fehlen jedoch detaillierte Informationen über den Leistungsfluss in einem Gebäude, um das Einsparpotential verschiedener Verbraucher erkennen und realisieren zu können. Einerseits können einzelne Stromzähler in Gebäuden keine Informationen darüber liefern, welcher Verbraucher wie viel Energie verbraucht. Andererseits ist der Einsatz von individuellen Zählern für jeden Verbraucher zu kosten- und wartungsintensiv.

Um die Leistungsaufnahme typischer Verbraucher in Büro-Umgebungen kosteneffizient und genau messen zu können, wurde im Rahmen des EnerFlow-Projekts ein softwarebasiertes Verfahren zur Leistungsmessung auf Ebene der Einzelverbraucher entwickelt. Hierzu werden Messdaten von bereits vorhandenen Sensoren genutzt, um durch gezielte Verarbeitung dieser Messdaten Regressionsmodelle zur Bestimmung des Leistungsbedarfs von Elektrogeräten zu erstellen. Die entwickelten Methoden sind nicht gerätespezifisch, sodass diese für beliebige Verbrauchertypen verwendet und in beliebigen Gebäuden eingesetzt werden können. Die einzige nötige Voraussetzung hierfür ist, Sensordaten mit Korrelation zur Leistungsaufnahme des Verbrauchers ermitteln zu können. Mit den im Rahmen von EnerFlow entwickelten Methoden ist es möglich, detaillierte Einblicke in die Ursachen von Energieverbräuchen in Bürogebäuden zu erhalten. Insbesondere die Gebäudebus-basierte Messung der Leistungsaufnahme für einzelne Schaltgruppen ermöglicht es, hohe Einsparpotentiale bei der Beleuchtung zu identifizieren. Somit bildet die in diesem Projekt entwickelte Methode die Grundlage, um detaillierte Einblicke in den Energieverbrauch von Gebäuden zu erhalten. Statt einer monatlichen oder jährlichen Abrechnung der verbrauchten Energie ermöglicht der Einsatz von Energiemodellen die Erstellung ausführlicher Berichte zum Energieverbrauch für jeden beliebigen Verbraucher, Zeitpunkt oder Zeitraum. Damit wird es möglich, die Auswirkungen von Energieeffizienzmaßnahmen ohne zusätzliche messtechnische Eingriffe zu bewerten. Da eine entsprechende Lösung mit geringem Aufwand in bestehenden Infrastrukturen nachgerüstet werden kann, ergibt sich eine hohe Anwendbarkeit.

Die mit Hilfe der software-basierten Energieverbrauchsmessung gewonnenen Einsichten bilden die Grundlage, um auf betriebswirtschaftlicher Basis gezielt rentable Energiesparpotentiale aufzuzeigen und umzusetzen. Somit legt diese Arbeit einen wichtigen Grundstein für ein energieeffizienteres Wirtschaften.

Projektpartner: Technische Universität Darmstadt und Deutsche Telekom AG

Projektlaufzeit: 01.03.2014 bis 31.08.2015

Mobiler-Campus-Charlottenburg (App): Plattform zur Untersuchung von QoE of Mobile Gaming

Justus Beyer

Gaming-QoE-Model

Mehr als die Hälfte aller Smartphone-Besitzer spielt täglich. Spiele stellen die am häufigsten genutzte App-Kategorie dar und machen einen erheblichen Teil des Gesamtumsatzes mit Apps aus. Dennoch ist bislang nur wenig über die vielfältigen Anwendungs-Kontexte, in denen gespielt wird, und die für die Qualitätswahrnehmung relevanten Faktoren bekannt. Was ist überhaupt „Qualität“ bei Spielen?

Wie kann man sie messen und wie wird sie von technischen Faktoren wie der Gerätegröße oder der Internetverbindung beeinflusst? Im Software Campus-Projekt „Gaming-QoE-Model“ geht es darum, Antworten auf diese Fragen zu finden. Dazu implementieren wir eine Software-Plattform, mit der Spiele über eine Netzwerkverbindung ohne lokale Installation bei dem Nutzer gestreamt werden können. Da die Ausführung des Spiels auf einem kontrollierten Server stattfindet, lassen sich experimentelle Parameter kontrollieren und deren Wirkung in Labor- und Feldversuchen mit Versuchspersonen erforschen.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.03.2013 bis 30.09.2015

Hauptspeicherdatenbanken in der Cloud

Tobias Mühlbauer

HDBC

Unternehmen und wissenschaftliche Projekte sehen sich mit der Herausforderung konfrontiert, dass zunehmend mehr Daten innerhalb immer kürzerer Zeitspannen in ihrer Datenbasis bearbeitet und aggregiert werden. Von diesen Daten unterliegt ein Großteil ständigen Änderungen, etwa durch die Geschäftsprozesse in Unternehmen oder durch neue Erkenntnisse und Messwerte in wissenschaftlichen Datensammlungen. Gleichzeitig wird es immer dringlicher, Entscheidungen schneller und genauer zu treffen. Diese Entscheidungen sind in den meisten Fällen von Analysen der zuvor genannten sich schnell ändernden Datenbasis abhängig. Es ist deshalb wünschenswert, analytische Anfragen zur Entscheidungsunterstützung auf dem aktuell gültigen Datenbestand effizient auswerten zu können. Eine zentrale Datenhaltung kann auf Grund der enormen Datenmengen und Anfragelast oft nicht hoch genug skaliert werden, um den Anforderungen gerecht zu werden. Um dieses Hindernis zu überwinden, müssen Daten auf mehrere physikalische Systeme verteilt werden. Ebenso scheitern zentralisierte Ansätze oft an der geographisch weltweiten Verteilung der Aktoren, welche mit einem Datenbestand interagieren. Auch hier kann eine Verteilung oder Replikation der Daten auf geographisch verteilten Systemen helfen, die Verfügbarkeit zu erhöhen und die Anfragegeschwindigkeit, nicht zuletzt durch geringere Latenzzeiten, zu optimieren.

Das Ziel des Forschungsprojekts war es, ein sich selbst organisierendes Informationssystem auf einer verteilten Infrastruktur zu konzipieren, welches den eingangs genannten Herausforderungen mit der Geschwindigkeit von Hauptspeicher-Datenbanken zu begegnen versucht. Die Knoten in dieser verteilten Infrastruktur betreiben dabei jeweils ein hochperformantes Hauptspeicher-Datenbanksystem. HyPer, ein am Lehrstuhl für Datenbanksysteme der Technischen Universität München entwickeltes modernes Hauptspeicher-Datenbanksystem, bietet sich hierbei als eine mögliche Datenbanklösung an; aber auch andere SQL-basierte Hauptspeicher-Datenbanksysteme sind einsetzbar. Diese Systeme ermöglichen sich schnell ändernde Datenbestände sowie gleichzeitig stattfindende effiziente Datenanalysen.

Projektpartner: Technische Universität München und Software AG

Projektlaufzeit: 01.06.2012 bis 31.05.2014

Organisation von Software-Entwicklungsartefakten zur Verbesserung von Wiederverwendung im Kontext der industriellen Software-Entwicklung

Veronika Bauer

IndRe

Wiederverwendung in der Software-Entwicklung und Wartung gilt als ein entscheidender Faktor, hohe Effizienz und Qualität zu erreichen und gleichzeitig Kosten und Zeitaufwand zu senken. Diese Attribute sind in Zeiten der globalisierten Software-Entwicklung für das Bestehen von Unternehmen höchst relevant. Gerade die ingenieurmäßigen Prinzipien des Software Engineering haben das Potential wie in anderen Industriezweigen zum Alleinstellungsmerkmal deutscher Software-Produzenten und -Dienstleister zu werden.

Es gibt bereits Unternehmen, in denen Wiederverwendung erfolgreich praktiziert wird. Gleichzeitig ist die Anzahl von Fällen, in denen Wiederverwendung nicht stattfindet oder die unerwünschten Effekte, beispielsweise hohe Klonraten, nach sich zieht, ungleich größer. Als Konsequenz stellt sich eine Vielzahl an Forschungsfragen, wie beispielsweise: Wie praktizieren Software-Unternehmen Wiederverwendung? In welchem Umfang und in welchem Kontext findet Wiederverwendung statt? Was sind Erfolgs- oder Risikofaktoren in Bezug auf erfolgreiche Wiederverwendung? Wie müssen, von Unternehmen, entwickelte Artefakte strukturiert sein, um eine erfolgreiche Wiederverwendung zu ermöglichen? Wie können erfolgsversprechende Strategien für Wiederverwendung in Unternehmen(sprozessen) integriert werden? Je bedeutender die Rolle von Software-Entwicklung und -Wartung in den Unternehmen ist, desto wichtiger ist es, auf diese Fragen zufriedenstellende Antworten zu finden.

Intelligente Informationsextraktion für das Entdecken und Verknüpfen von Wissen

Sebastian Krause

Intellektix

Heutige IT-Systeme sind mit einer zunehmenden Flut von Informationen verschiedener Ausprägung konfrontiert. Dazu zählen neben klassischen strukturierten Daten in Form von Datenbanken vor allem Dokumente, d.h. Informationen in natürlicher Sprache. Sollen IT-Systeme nun einen Nutzer in die Lage versetzen, Daten aus allen möglichen Informationsquellen gleichartig finden und nutzen zu können, so muss ein solches System auch in der Lage sein, menschliche Sprache verarbeiten zu können. Sprachtechnologie als Werkzeug zur Erstellung solcher Systeme hat in den letzten Jahren viele Fortschritte gemacht, die auch einem breiten Publikum erkennbar waren. Ein Beispiel hierfür ist „Siri“, welches ein natürlchsprachiges, verbales Interface zur Interaktion mit einem Smartphone bereitstellt. IBMs „Watson“, ein Frage-Antwort-System, welches in der Quiz-Fernsehsendung „Jeopardy!“ gegen Menschen antrat, war zusätzlich zum Annehmen von Anweisungen auch in der Lage, die Bearbeitung der Anfrage auf unstrukturierten Daten durchzuführen. Googles „Knowledge Graph“ ist eine Wissensdatenbank, welche eingesetzt wird, um die Vieldeutigkeit der menschlichen Sprache sowie Ungenauigkeit durch verkürzte Ausdrucksweise aufzuschlüsseln und aufzulösen.

Ziel des Vorhabens ist es, einen Prototyp eines Informationsextraktionssystems zu entwickeln, welcher es erlaubt, automatisch aus Texten bestimmter Zieldomänen strukturierte Daten zu entnehmen. Das System soll hierbei mittels vorhandenen Domänenwissens, zum Beispiel in Form von unternehmensinternen Datenbanken oder im Internet frei verfügbaren Daten, trainiert werden, ohne dass manuelles Eingreifen nötig ist. Der Schwerpunkt der Arbeit wird sein, sowohl Abdeckung als auch Genauigkeit des Systems gegenüber dem aktuellen Technikstand zu verbessern. Die aus diesen Vorhaben resultierende Technologie liefert einen Teil der nötigen Bausteine für fortgeschrittene Informationssysteme, die menschliche Nutzer intelligent bei der Befriedigung ihres Informationsbedürfnisses unterstützen. Außerdem bilden die entstehenden Technologien wichtige Bestandteile beispielsweise für Informationssysteme, die Medien permanent nach neuen Berichten oder Meinungen zu einem Thema oder einer Entität durchsuchen und dadurch regelmäßige Überblicke über die Medienpräsenz des Themas oder der Entität erstellen sowie Trends in der Berichterstattung und öffentlichen Meinung zu einem Thema aufzeigen können.

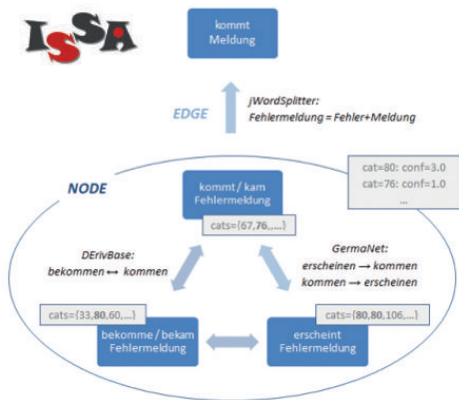
Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und Scheer Group GmbH

Projektlaufzeit: 01.03.2013 bis 28.02.2014

Inferenzbasierte Suche auf Supportanfragen

Kathrin Eichler

ISSA



Ziel des Projektes ISSA war die Entwicklung einer auf textueller Inferenz basierenden Suche, die eingesetzt werden kann, um für eine eingehende Anfrage relevante Dokumente zu finden. Anders als bei herkömmlichen Suchfunktionen sollten hierfür der Anfragetext und die bereits im System erfassten Dokumente auf inhaltlicher Ebene analysiert und miteinander verglichen werden.

Innerhalb des Projektes wurde hierfür ein Ansatz entwickelt, der die bestehenden Dokumente in einen Entailment-Graphen umwandelt und diesen in den Suchprozess einbindet. Auf der Grundlage einer durchgeführten Datenanalyse wurden für die Erstellung des Graphen externe Sprachressourcen sowie Tools zur linguistischen Vorverarbeitung verwendet. Der Fokus des Projektes lag auf der Verarbeitung deutschsprachiger Daten. Zum Einsatz kamen daher unter anderem das Derivationslexikon DERivBase, das lexikalisch-semantische Netz GermaNet, der jWordSplitter zur Zerlegung von Komposita sowie der MaltParser zur Erkennung von syntaktischen Abhängigkeiten.

Die Auswertung des entwickelten prototypischen Systems erfolgte auf einem Datensatz aus Supportanfragen, für die mit Hilfe des Systems passende Kategorien ermittelt wurden. Unter anderem wurde der Einfluss der eingesetzten Tools und Ressourcen auf das Gesamtergebnis ausgewertet. Hier stellte sich heraus, dass sich insbesondere die Verwendung eines Derivationslexikons sowie die Zerlegung von Komposita positiv auf das Ergebnis auswirken. Insgesamt konnte mit der eingesetzten Technologie eine Verbesserung der Kategorisierungsergebnisse erzielt werden.

Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und Software AG

Projektlaufzeit: 01.03.2013 bis 28.02.2015

Gitterbasierte Kryptographie für die Zukunft

Rachid El Bansarkhani

IT-GiKo

Nahezu alle modernen Unternehmensanwendungen werden mit Hilfe von Verfahren der Kryptographie abgesichert, um Schutzziele wie Vertraulichkeit, Authentizität, Integrität, Privatheit und Zurechenbarkeit der Daten zu garantieren. Dazu zählen beispielsweise Verschlüsselungs-, Signatur-, Identifikations- und Hashverfahren, die zum Standardwerkzeug von IT-Sicherheitslösungen zählen. Mit Hilfe von Quantencomputern ist es theoretisch möglich, effizient zu faktorisieren und diskrete Logarithmen zu berechnen, wodurch die in der Praxis eingesetzten klassischen Verfahren wie RSA, ECC und DSA gebrochen werden können und damit keinen Schutz mehr garantieren.



Sogenannte gitterbasierte Kryptoverfahren, deren Resistenz gegen Quantencomputer sehr stark angenommen wird, werden als alternative Kandidaten betrachtet. Sie übertreffen die klassischen Verfahren in der Regel sogar im Hinblick auf die Performance. Um zukünftige Anwendungen und IT-Infrastrukturen mit gitterbasierten Sicherheitslösungen auszustatten, ist es erforderlich, die Praktikabilität gitterbasierter Verfahren früh genug zu erforschen. Im Projekt „IT-GiKo“ wurde die Praktikabilität von gitterbasierten Verfahren weiter vorangetrieben. Es wurden verschiedene gitterbasierte Verfahren betrachtet und implementiert, um deren Effizienz und Sicherheitsmerkmale im Hinblick auf den Einsatz in zukünftigen Anwendungen zu beurteilen und zu messen. Darüber hinaus beinhaltet dieses Projekt die Verbesserung und Entwicklung neuer gitterbasierter Sicherheitstechnologien, wie z.B. intelligente Kompressionstechnologien zur Gestaltung von effizienten Signaturverfahren oder die Entwicklung von kompakten Verschlüsselungsalgorithmen für große Nachrichtenmengen.

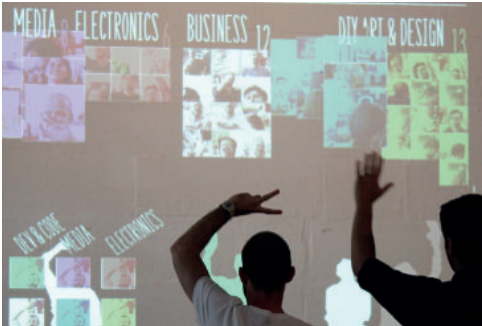
Projektpartner: Technische Universität Darmstadt und SAP SE

Projektlaufzeit: 01.01.2013 bis 31.12.2014

Immediate Usability für Interaktive Public Displays

Robert Walter

IUPD



Die Verbreitung von Displays (z.B. als Werbefläche) im öffentlichen Raum nimmt mehr und mehr zu. Durch Touchtechnologien und Gestenerkennung werden diese Bildschirme interaktiv. Um die Konversionsrate der Bildschirme zu maximieren, müssen die Bildschirme 1. dem Passanten kommunizieren, dass sie interaktiv sind, 2. die Interaktion erklären und 3. den Passanten zur Interaktion motivieren. In diesem Projekt werden Ansätze zur Lösung dieser drei Probleme erarbeitet.

Der Fokus des Projekts liegt dabei bei der sensorisch gestützten Interaktion mit großformatigen Displays, zum Beispiel über Kameras oder spezielle Tiefensensoren. Dadurch ergeben sich aus Sicht der Entwickler und Anbieter interaktiver öffentlicher Displays enorme Möglichkeiten: Der gesamte Körper des Benutzers kann zur Eingabe dienen. Damit ist die Interaktion auch über die Entfernung möglich und kann somit sehr früh und subtiler erfolgen. Der Passant interagiert bereits im Vorbeigehen, was mit Touchscreen-Lösungen nicht möglich ist. Der Übergang vom Passanten zum Benutzer wird somit kontinuierlich. Ziel ist es, die Interaktion mit öffentlichen Displays so zu gestalten, dass sie für die Passanten interessant, greifbar und vor allem verständlich ist.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.04.2013 bis 31.03.2015

Verfeinerte Extraktion von Best-Practices aus Software Repositories durch Kombination von automatisierten Verfahren mit Expertenwissen

Sebastian Proksch

KaVE

Entwickler verwenden jeden Tag Frameworks. Jedes neue Framework stellt dabei den Entwickler vor einen hohen Einarbeitungsaufwand. Aber auch für die wiederkehrende Benutzung müssen sich Entwickler Best-Practices merken, was bei der schier unendlichen Anzahl der zu Verfügung stehenden Technologien oft eine schwere Aufgabe ist und viel Übung erfordert.

Ein Lösungsansatz ist die Unterstützung der Entwickler bei der Arbeit in ihrer Entwicklungsumgebung durch Tools, die ihnen Hilfestellungen bieten. Es wurde bereits mehrfach gezeigt, dass solche Systeme realisierbar sind. Die bereitstehenden Tools lassen sich dabei in zwei Gruppen einordnen: (1) Ein klassischer Ansatz ist das Bereitstellen von Regelwerken oder Codebeispielen, die durch Experten aufwändig erzeugt und gepflegt werden. (2) In neueren Ansätzen werden tool-gestützt große Mengen an Beispielen analysiert und dabei automatisiert erlernt, wie Entwicklern geholfen werden kann. Beide Ansätze haben für sich genommen enorme Nachteile: Bei ersterem ist durch die beschränkte Leistungsfähigkeit eines Experten meist der Umfang sehr gering und expertengepflegte Tools sind für Unternehmen sehr teuer in Erstellung und Wartung. Bei zweitem besteht das Wissen meist aus einer einfachen statistischen Auswertung. Dies führt jedoch dazu, dass auf Grund statistischer Relevanz auch unerwünschte Anti-Patterns gelernt werden. Innerhalb des Vorhabens KaVE wird in einem neuartigen Ansatz untersucht, ob beide Ansätze kombiniert werden können. Hierdurch sollen die Vorteile vereint und zeitgleich die Nachteile vermieden werden. Experten sollen in den automatisierten Lernprozess integriert werden und direkten Einfluss auf das Ergebnis nehmen. So wird es möglich die Fehlerrate automatisierter Verfahren zu reduzieren und zugleich die Belastung der Experten zu verringern. Die Arbeit wird schneller und mit höherer Qualität erledigt.

Der Einsatz der resultierenden Tools verspricht eine Effizienzsteigerung der Entwickler, diese wurde jedoch bisher noch nie im Praxiseinsatz untersucht. Durch die enge Kooperation mit einem Industriepartner ist es uns im Rahmen des Vorhabens möglich, diese Tools in einem realen Entwicklungsumfeld und mit einer großen Anzahl an professionellen Softwareentwicklern zu evaluieren. Die Qualität der Erkenntnisse einer solch praxisnahen Evaluation liegt weit über dem rein theoretischer Untersuchungen, wie sie in bisherigen Studien üblich sind. Daher ist die ausführliche Evaluation ein Schwerpunkt des Vorhabens. Im Rahmen des Projekts wird außerdem untersucht, wie einmal durchgeführte Nutzerstudien wiederverwendbar konserviert werden können, um eine zukünftige Simulation des Verhaltens der Entwickler zu ermöglichen. Dies würde die praxisnahe Evaluation eines neuen Tools ohne erneute Nutzerstudie erlauben.

Projektpartner: Technische Universität Darmstadt und DATEV eG

Projektlaufzeit: 01.01.2013 bis 30.04.2015

Kosteneffiziente Bereitstellung Cloud-basierter Unternehmensanwendungen unter Berücksichtigung von Dienstgüteanforderungen

Ronny Hans

KoBeDie

Seit Jahren ist Cloud Computing ein Bereitstellungsmodell für IT-Dienstleistungen, welches stetig an Bedeutung gewinnt. Während anfangs Infrastrukturdienstleistungen den Schwerpunkt bildeten, gewinnen multimediale Anwendungen zunehmend an Bedeutung. Diese zeichnen sich insbesondere durch hohe Anforderungen an die Dienstgüte, z. B. einer geringen Latenz, aus. Bisherige Untersuchungen haben gezeigt, dass durch die weitestgehend zentralisierte Cloud-Infrastruktur nur ein Teil der potenziellen Konsumenten mit solchen Diensten versorgt werden kann. Wollen Anbieter von Cloud-Infrastrukturdiensten an dem wachsenden Markt multimedialer Dienstleistungen partizipieren, so ergeben sich verschiedene Herausforderungen. Um die Latenz zwischen Anbieter und Konsumenten zu reduzieren, müssen Cloud-Ressourcen in der Nähe potenzieller Konsumenten platziert werden, was eine effiziente Standortplanung beim Ausbau einer dezentralen Cloud-Infrastruktur bedarf. Zum anderen muss die vorhandene Infrastruktur effizient genutzt werden, um eine größtmögliche Anzahl von Konsumenten mit multimedialen Diensten zu versorgen. Dies erfordert, dass dynamisch auf die Nachfrage reagiert und entsprechend Ressourcen zugewiesen werden können.

Den Schwerpunkt dieses Forschungsvorhabens bildet die Entwicklung von Optimierungsalgorithmen zur Rechenzentrenausswahl, die es Diensteanbietern ermöglicht, durch eine intelligente Auswahl von Ressourcen, eine möglichst große Anzahl von Konsumenten kostenminimal mit multimedialen Diensten zu versorgen. Da trotz der Komplexität eines solchen Zuordnungsproblems Lösungen in kürzester Zeit ermittelt werden müssen, gilt es Verfahren zu konzipieren, die abhängig vom Anwendungsfall, möglichst genaue Ergebnisse mit geringem Rechenaufwand ermitteln. Hinsichtlich des Nutzerverhaltens werden dabei zwei unterschiedliche Annahmen zu Grunde gelegt. In einem ersten Szenario wird davon ausgegangen, dass die Nachfrage nach Cloud-Diensten und deren Dienstgüteanforderungen a priori bekannt ist. Für Planungsszenarien, die in der Zukunft liegen, unterliegen die genannten Größen aber Unsicherheiten und können nicht mehr genau vorhergesagt werden. Deshalb sollen im Rahmen dieses Forschungsvorhabens auch Verfahren entwickelt werden, die in der Lage sind, Zufallsgrößen bei der Ressourcenallokation zu berücksichtigen, um mögliche Zukunftsszenarien zu untersuchen.

Kontextsensitivität zur Steigerung der Sicherheit und User Experience mobiler Anwendungen

Christian Jung

KoSIUX

Kontextsensitivität bezeichnet das Anpassen des Systemverhaltens an den Anwendungskontext basierend auf Informationen zur Nutzungssituation, wie Umgebungsinformationen, Zustand und Verfassung des Anwenders oder den aktuellen Ort der Nutzung. Damit können Sicherheitsrichtlinien flexibel und adäquat zur jeweiligen Situation des Nutzers und des Geräts angepasst werden. Weiterhin können Apps dynamisch auf die Nutzungssituation reagieren, um das Nutzungserlebnis des Anwenders (engl. User Experience) zu verbessern.

Besonders im Bereich mobiler Endgeräte birgt Kontextsensitivität ein großes Potenzial zur Steigerung der Sicherheit und User Experience. Mobile Geräte werden in unterschiedlichsten Nutzungssituationen eingesetzt, da Anwender im Berufs- und Privatleben häufig zwischen verschiedenen Aufgaben und Standorten wechseln. Zusätzlich ist ein stetiger Wandel hin zur Überlappung von dienstlicher und privater Nutzung zu verzeichnen. Durch die automatische Erkennung von Kontexten und Kontextwechseln kann der Nutzer bei der Trennung verschiedener beruflicher und privater Arbeitsabläufe unterstützt werden, und das Gerät kann sich dem jeweiligen Anwendungsfall anpassen. Ein Beispiel hierfür ist das kontrollierte Verhindern des Zugriffs auf geschäftliche Daten im privaten Umfeld.

Ziel des Projekts KoSiUX war es, ein besseres Verständnis von Kontexten und deren Spezifikation im Bereich von Sicherheitsrichtlinien zu ermöglichen und mobile Endgeräte um Mechanismen zur einfachen Nutzung von Kontexten zu erweitern. Hierzu wurde ein Modell zur Beschreibung von Kontexten entwickelt, welches die Aussagekraft einer Kontextinformation und deren Verlässlichkeit mit in die Kontextentscheidung einfließen lässt. Weiterhin wurde eine Methode zur teilautomatischen Erhebung und Modellierung von Kontextbeschreibungen entwickelt, welche eine wiederholbare und vergleichbare Spezifikation von Nutzungssituationen ermöglicht. Hierbei kann die Erstellung der Kontextbeschreibungen hinsichtlich Trefferquote (engl. recall) und Genauigkeit (engl. precision) optimiert werden. Abschließend wurde ein Prototyp für Android entwickelt, welcher mit Hilfe der generierten Kontextbeschreibungen konfiguriert wird und anderen Anwendungen Kontextinformationen zur Verfügung stellt. Die Komponente wurde erfolgreich in das Sicherheitsframework IND²UCE (Integrated Distributed Data Usage Control Enforcement) integriert und erlaubt dort kontextabhängige Sicherheitsentscheidungen.

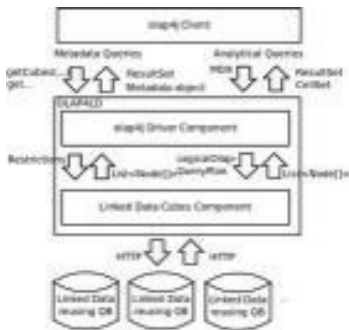
Projektpartner: Fraunhofer-Verbund IuK-Technologie und DATEV eG

Projektlaufzeit: 01.01.2013 bis 30.04.2015

Linked Data Cubes - Integration und Analyse Verteilter Statistischer Datensätze durch Linked Data

Benedikt Kämpgen

LD-Cubes



Das Projekt LD-Cubes (<http://www.linked-data-cubes.org/>) präsentiert ein Framework zur Entwicklung von Analyseapplikationen über verteilte Statistikdaten.

Applikationen (Abbildung 1) integrieren semi-automatisch alle verfügbaren, auch heterogenen Datensätze zu einem Globalen Datenwürfel und erlauben Online Analytical Processing (OLAP, oberer Teil der Abbildung). OLAP hat sich als intuitive Analyseoberfläche für Entscheidungsträger bewährt. Den Applikationen sind jegliche Datensätze verfügbar, die im Intranet oder Web mittels weit-verbreiteter Best Practices (Webprotokoll HTTP, Graphdatenmodell RDF, Semantic Web Vokabular RDF Data Cube) veröffentlicht sind (unterer Teil der Abbildung).

Kernkomponente ist die Open-Source OLAP Engine OLAP4LD. Sie lädt automatisch erforderliche Datensätze, integriert Daten auf Basis flexibel hinzufügbarer Hintergrundinformationen und führt effizient Abfragen aus. Beispielsweise wurde eine Demo (Linked Data Cubes Explorer, <http://ldcx.linked-data-cubes.org/projects/ldcx/>) entwickelt, die z.B. Bürgern erlaubt, die Wachstumsrate des Bruttoinlandsprodukt von Eurostat mit Ergebnissen einer Umfrage über Arbeitslosigkeit des GESIS Instituts zu vergleichen und Trends abzulesen.

Zusätzliche Experimente zeigen, dass auch eine kollaborative Analyse von Statistikdaten mittels Semantic MediaWiki sowie die automatische Ableitung von neuen Datensätzen durch Hintergrundinformationen, z.B. Umrechnungskursen, möglich sind. Im Anschluss an das Projekt ist es geplant, OLAP4LD-Applikationen zur Analyse des Globalen Datenwürfels in der Industrie, z.B. beim Partner SAP, einzusetzen.

Projektpartner: Karlsruher Institut für Technologie (KIT) und SAP SE

Projektlaufzeit: 01.01.2013 bis 31.12.2014

Domänen-Unabhängiges Modulares Scheduler-Framework für Heterogene Multi- und Many-Core Architekturen

Anselm Busse

ModSched

Rechnersysteme haben in unserer technisierten Welt einen sehr hohen Durchdringungsgrad erreicht und sind in fast jedem elektronischen Gerät zu finden – von einer einfachen Waschmaschine bis hin zur komplexen Fertigungsanlage. Fast all diesen Systemen ist gemein, dass sie eine Art Betriebssystem benötigen, welches die Aufgaben des jeweiligen Systems koordiniert und Ressourcen verwaltet. Eine zentrale Fragestellung innerhalb eines Betriebssystems ist dabei, welche Aufgabe zu welchem Zeitpunkt zu bearbeiten ist. Diese Aufgabe übernimmt der Scheduler. Dabei sind verschiedene Betriebsziele zu beachten. Bei der Steuerungssoftware eines Verbrennungsmotors steht z.B. die Einhaltung von Fristen im Vordergrund, um die Funktionstüchtigkeit des Motors zu gewährleisten, während bei einem Desktoprechner die Antwortzeit des Systems von stärkerer Bedeutung ist, um eine gute Nutzererfahrung sicher zu stellen. Der Scheduler wird dabei für gewöhnlich systemspezifisch entworfen und implementiert. Diese Entwicklung ist meist sehr aufwendig und kostenintensiv.

Die Forschungsergebnisse des präsentierten Projekts ermöglichen den Entwicklungsaufwand im Bereich des Schedulers zu verringern, die Entwicklungszeiten zu verkürzen und somit schnellere Innovationszyklen. Hierfür wurde ein Domänen-Unabhängiges modulares Scheduler-Framework entwickelt, dass im Rahmen des Projekts prototypisch sowohl für Linux als auch FreeBSD implementiert wurde. Das Framework ermöglicht die Implementierung von beliebigen Scheduling-Algorithmen. Durch das Framework wird anders als bisher kein umfangreiches Vorwissen über den Aufbau und die Struktur des umgebenden Betriebssystems benötigt.

Darüber hinaus ist es möglich, im Gegensatz zum aktuellen Stand der Technik, eine Scheduler-Implementierung oder modularisierte Teile derer ohne Portierungsaufwand für mehrere Betriebssysteme zu nutzen. Dieses Vorgehen verringert sowohl Entwicklungs- als auch Verifikationskosten und ermöglicht eine schnelle Anpassung des Schedulers an neue Betriebsziele und Anforderungen.

Projektpartner: Technische Universität Berlin und Robert Bosch GmbH

Projektlaufzeit: 01.04.2013 bis 31.12.2014

Modulares Vertrauensmanagement für dynamische Dienste

Sascha Hauke

Move4Dynamic

Vertrauen und Vertrauenswürdigkeit sind wichtige Treiber von sozialem und wirtschaftlichen Fortschritt. Dies gilt auch für den digitalen Fortschritt, etwa im Hinblick auf das derzeitige und zukünftige Internet. Klassische Verfahren zur Schaffung von Vertrauen sind in digitalen Räumen jedoch nur bedingt anzuwenden. Im Projekt MoVe4Dynamic wurden deshalb Verfahren entwickelt, die es ermöglichen, die Vertrauenswürdigkeit eines digitalen Gegenübers zu berechnen und somit einem Nutzer eine fundierte Grundlage zur Bildung von Vertrauen zur Verfügung zu stellen.

Hierzu wurde zunächst als konzeptionelle Grundlage die Architektur eines modularen Vertrauensmanagementsystems vorgestellt, die es ermöglicht aus verschiedenen Komponenten zur Vertrauensberechnung, -darstellung und Entscheidungsfindung ein Vertrauensmanagementsystem für einen bestimmten Anwendungsfall konkret zusammenzustellen. Diese wurde in einem Softwareprototyp realisiert.

Der Fokus der weiteren Arbeit lag in der Konzeption akkurater Vorhersagemethoden und der Weiterentwicklung von Schätzern für die Vertrauenswürdigkeitsvorhersage. Neben der Betrachtung statistischer Verfahren, aus denen eine Reihe neuartiger Schätzer abgeleitet werden konnten, wurden auch sogenannte überwachte Lernverfahren auf ihre Eignung getestet, unter anderem durch Validierung mit Daten, die aus real existierenden Systemen gewonnen wurden.

Mathematische Modelle zur Vertrauensrepräsentation wurden weiterentwickelt, sowohl hinsichtlich der darstellbaren Granularität von vertrauensrelevanten Daten als Datengrundlage, als auch bezüglich der Kombination von vertrauensrelevanten Daten.

Weiterhin wurde ein Verfahren zur Visualisierung von multikategorialen Vertrauensbewertungen entwickelt. Dieses ermöglicht es in einem Diagramm Informationen zu Vertrauenswürdigkeit eines Gegenübers anhand verschiedener Aspekte darzustellen. Eine durchgeführte Nutzerstudie hat gezeigt, dass dieses Darstellungsverfahren intuitiv ist und angenommen wird.

Die im Projekt MoVe4Dynamic erzielten Fortschritte sind wissenschaftlich anschlussfähig und stellen die Basis für weitere wissenschaftliche Arbeiten am Lehrstuhl und in Kooperationen mit anderen (internationalen) Hochschulgruppen dar.

Projektpartner: Technische Universität Darmstadt und SAP SE

Projektlaufzeit: 01.01.2013 bis 31.12.2014

Dezentrale Steuerung von cyber-physischen Produktionssystemen

Julius Pfrommer

OptimusPlant

Produzierende Unternehmen sind heute mit dem Problem steigender Komplexität von Produkten bei immer kürzeren Lebenszyklen konfrontiert. Gleichzeitig haben in der globalisierten Wirtschaft die Marktunsicherheiten auf Kundenseite und bei Lieferanten zugenommen. So beträgt etwa der Produktionszeitraum für ein neues Handymodell weniger als ein halbes Jahr. Und die Nachfrage kann, abhängig von den Modellen der Konkurrenzhersteller und vorteilhaften Testberichten, kaum vorhergesagt werden. In der Folge müssen die Produktionsanlagen immer häufiger umgebaut und die zugehörige Steuerung angepasst werden. Für voll automatisierte und verkettete Anlagen kann das u.U. Monate dauern. Im Fall der Textil- und Elektronikbranche wird die erforderliche Flexibilität durch den weiterhin hohen Anteil menschlicher Arbeitskraft in der Fertigung erreicht. Dies hatte jedoch eine Abwanderung weiter Teile der Industrien in Niedriglohnländer zu Folge.

Die digitale Steuerung von industriellen Prozessen und die Verkettung von Anlagen und Maschinen über Feldbusse sind nun seit knapp 40 Jahren Stand der Technik und der Kern vieler „cyber-physischen Systeme“. Und bis zuletzt haben auf dem Gebiet durch die fortschreitende Entwicklung der Technologie und die Standardisierung Verbesserungen stattgefunden. Im Zuge der Entwicklungen rund um „Industrie 4.0“ soll nun durch die Verbindung der Automatisierungstechnik mit Internet-Technologien eine neue Klasse von Anwendungen und Geschäftsmodellen entstehen. Gerade die Flexibilisierung und Wandlungsfähigkeit von automatisierten Anlagen bekommt dabei viel Aufmerksamkeit. Die Automatisierungstechnik soll in Zukunft nicht nur denselben Ablauf (mit vorherbestimmten Variationen) endlos wiederholen, sondern zur Laufzeit flexibel auf immer neue Situationen reagieren können. Es ist aber noch unklar, wie die zukünftige Steuerungsarchitektur von Produktionsanlagen aussehen muss, um diese Anforderungen zu erfüllen.

Im Software Campus-Projekt OptimusPlant soll untersucht werden, wie die Planung und Steuerung in cyber-physischen Produktionssystem (CPPS) nicht zentral vorgegeben, sondern dezentral verteilt werden kann. Das Ergebnis des Projekts ist zum einen die Konzeptionierung und prototypische Umsetzung einer Softwareplattform, über die Produktionsanlagen mit einem service-orientierten Ansatz betrieben werden können. Zum anderen werden auf Basis der Softwareplattform ausgewählte zentrale und dezentrale Steuerungsansätze umgesetzt und vergleichbar gemacht.

Projektpartner: Karlsruher Institut für Technologie (KIT) und Robert Bosch GmbH

Projektlaufzeit: 01.03.2014 bis 29.02.2016

Plattform zur effizienten Analyse und sicheren Komposition von Softwarekomponenten

Ben Hermann

PEAKS

Softwaresysteme nehmen eine immer zentralere Rolle in der Abwicklung des täglichen Lebens ein. Informationstechnik dringt teilweise unbemerkt in jeden Aspekt des Alltags ein. Dabei steigt nicht nur die Menge der gesammelten Daten, sondern auch die Zahl an Softwaresystemen. All diese Softwaresysteme müssen so beschaffen sein, dass sie verantwortungsvoll mit den Daten und Berechtigungen ihrer Nutzer umgehen. Bisherige Verfahren verlassen sich bei der Prüfung von Softwaresystemen auf ihre Sicherheit noch oftmals stark auf manuelle Prüfungen des Quellcodes – sogenannte Code Reviews. Dies ist für eine steigende Anzahl von zu prüfenden Systemen und einer zunehmenden Komplexität nicht effizient.

Ein fundamentaler Baustein in der Verbesserung der Effizienz moderner Softwareentwicklung ist die Wiederverwendung von Komponenten zur Komposition von Anwendungen. Damit schreibt ein Entwickler nicht jeden Teil seiner Applikation selbst, sondern verlässt sich u.a. auf das Betriebssystem, die Programmierplattform und weitere Bibliotheken. Diese wiederverwendeten Komponenten besitzen jedoch in dem Sicherheitskontext einer Anwendung unter Umständen mehr Berechtigungen als sie benötigen.

PEAKS macht es Softwareentwicklern einfach, Komponenten zu finden, die nur die Systemressourcen nutzen, die sie wirklich brauchen. Dadurch liefert das Projekt einen wichtigen Beitrag in der Vermeidung von Sicherheitslücken, die durch exzessive Berechtigungsanhäufung entstehen. Durch theoretisch fundierte, effiziente Codeanalysen können diese Untersuchungen auch kostengünstig bei jedem Versionsupdate durchgeführt werden.

Policy-getriebene Konextualisierung in ereignisbasierter Middleware

Tobias Freudenreich

Poker

Wir leben in einer Zeit der zunehmenden Digitalisierung, was durch Trends wie das Internet der Dinge verdeutlicht wird. Speziell in der Produktion und den angrenzenden Bereichen der Wertschöpfungskette werden reaktive Computersysteme zur Steuerung und Optimierung von Prozessen genutzt. Als Schlagwort hat sich in den letzten Jahren Industrie 4.0 durchgesetzt. Dabei werden vielerlei Sensordaten in Prozessen genutzt, um diese zu optimieren oder zu steuern.

Die große Herausforderung für die Softwaresysteme besteht darin, dass das „mehr“ an Daten effizient verarbeitet werden muss, um von Nutzen zu sein. Dies hat den Begriff Big Data geprägt. Big Data umfasst die sogenannten drei Vs: Volume, Velocity und Variety. Damit ist gemeint, dass sowohl die schiere Menge an zu verarbeitenden Daten sehr groß geworden ist (Volume), als auch deren Entstehungsfrequenz und die damit verbundene zeitnahe Verarbeitung (Velocity), sowie die Heterogenität (Variety) zugenommen hat.

Dadurch werden die entsprechenden Softwaresysteme sehr komplex und erfordern zum Gebrauch hoch qualifizierte Informatiker. Wünschenswert wäre es allerdings, dass die Möglichkeiten dieser Systeme auch von Domainexperten genutzt werden können, die sich eben nicht durch ein tiefes Informatikwissen, sondern durch großes Fachwissen auszeichnen.

Ein ähnliches Ziel verfolgt auch SQL bei Datenbanken: eine deklarative Sprache zum einfachen Umgang mit relationalen Daten zur Verfügung zu stellen. POKER erlaubt es Endanwendern komplexe Situationen und die Reaktionen darauf auf einfache Weise in einer deklarativen Sprache zu formulieren. Eine Middleware setzt diese Beschreibungen dann automatisch um. Die Middleware kann dabei von Computerexperten auf die Bedürfnisse des Unternehmens angepasst werden (ähnlich zu Datenbankexperten, die z.B. ein Datenbankschema erstellen). Zum Beispiel kann ein Logistikbeauftragter folgende Anweisung formulieren: „Falls zwei Gefahrgüter zu nah beieinander sind, informiere die Zentrale“. Die Middleware ermittelt dann welche Sensoren und statische Datenquellen dazu gebraucht werden und generiert und managt automatisch Codes zur Überprüfung der Anweisung.

Projektpartner: Technische Universität Darmstadt und Software AG

Projektlaufzeit: 01.01.2013 bis 31.03.2015

Optimierung von gamifizierten Anwendungen durch Echtzeitvorhersagen von Nutzerverhalten und Anreizen

Stefan Tomov

PreTIGA



Gamifizierung ist ein neuer Trend zur Steigerung der Anreize und des Engagements von Personen durch den Einsatz von Spieldesign-Elementen in einem Kontext, der nichts mit Spielen zu tun hat. Dadurch werden viele neuartige und kreative Anwendungen möglich. Bei Gamifizierung steht immer im Vordergrund, durch den Kontext vorgegebene Probleme zu lösen. Zum Beispiel findet Gamifizierung in eLearning oder ökologischen Systemen Anwendung, um Herausforderungen anzugehen, welche einen hohen Motivationsgrad von Teilnehmern solcher Systeme fordern. Zwei typische Beispiele sind die Steigerung der Effizienz, Effektivität und somit der erfolgreichen Kursabschlüsse in Online-Lernplattformen, oder die Steigerung der Energieeffizienz durch Anreiz zum bewussten Umgang mit Energie im Eigenheim. Der Erfolg von Spieldesign-Techniken hängt dabei primär von dem Anreizlevel der Spieler und ihrem Verhalten im Spiel ab. Ansätze, welche ein großes Potenzial zur Optimierung der Anreizsetzung liefern, sind eine feine Granularität des Spieldesigns und Personalisierung des Spielerlebnisses.

Ein erster Schritt hin zur Echtzeitoptimierung des Spielerverhaltens in gamifizierten Systemen ist die nicht-invasive Echtzeitvorhersage von Zuständen, die das Erreichen der Ziele des Systems behindern. Ein Beispiel hierfür ist das frühzeitige Verlassen eines Online-Kurses. Zur Verhinderung dieser Zustände und zur Optimierung von Spielumgebungen werden im Projekt PreTIGA proaktive Personalisierungstechniken zur Anreizoptimierung entwickelt und evaluiert. Dabei werden der Spielertyp und das individuelle Verhalten eines Spielers berücksichtigt, um sein Spielerlebnis zu verbessern.

Induktive Entwicklung von Referenzmodellen mithilfe von Maschinensemantik und deren praktische Evaluation

Jana Rehse

RefMod@RunTime

Die Unternehmensmodellierung beschäftigt sich mit der Abbildung von Ablauf- und Aufbauorganisation, also den Prozessen und Strukturen eines Unternehmens. Diese Modelle werden hauptsächlich zu Analyse- und Dokumentationszwecken verwendet, beispielsweise bei der Prozessoptimierung. Insbesondere bei der Einführung und dem Management von Informationssystemen sind diese Modelle für die meisten Unternehmen unverzichtbar. Nach wie vor werden Prozessmodelle häufig manuell erstellt. Dies ist auf der einen Seite aufwändig und fehleranfällig, andererseits sehr subjektiv, da sowohl die Anforderungen als auch das finale Modell vom persönlichen Hintergrund und der Zielsetzung des Modellierers abhängig sind.

Um Zeit und Kosten zu sparen, können bereits existierende Modelle verwendet werden, um neue Modelle zu erstellen. Diese sogenannten Referenzmodelle bilden gängige Praktiken in der jeweiligen Branche ab und können angepasst werden, um die individuellen Anforderungen von Unternehmen zu erfüllen. So profitiert ein Unternehmen von bekannten „best practices“ und spart gleichzeitig Zeit und Aufwand.

Es existieren bereits Referenzmodelle für einige Branchen oder Geschäftsbereiche, aber die meisten Unternehmen können nicht von den Vorteilen der Referenzmodellierung profitieren, da es keine Referenzmodelle von ausreichend hoher Qualität gibt. Es besteht also ein hoher Bedarf zur Entwicklung neuer Referenzmodelle. Diese können entweder deduktiv, also auf der Grundlage allgemeiner Theorien und Prinzipien, oder induktiv, also durch die Analyse individueller Unternehmensmodelle, entwickelt werden.

Ziel des Projektes RefMod@RunTime ist es, eine neue Methode zur induktiven Entwicklung von Referenzmodellen zu erarbeiten. Dazu soll, anders als in den bestehenden Ansätzen, die Ausführungssemantik der Ausgangsmodelle beachtet werden. In jedem Individualmodell werden mögliche Ausführungen identifiziert. Diese werden verglichen, um Gemeinsamkeiten in der Modellmenge herauszuarbeiten. Mit Methoden des Process Mining wird aus diesen Sequenzen ein neues Modell generiert, was als vorläufiges Referenzmodell verwendet werden kann. Im zweiten Teil des Projektes werden in Zusammenarbeit mit dem Industriepartner praktische Erfahrungen mit dieser Methode gesammelt. Einerseits geht es dabei darum, herauszufinden, ob induktive für den Anwender interessant und praktikabel sind und wie hoch ihre Akzeptanz ist. Andererseits sollen gemeinsam Anwendungsszenarien im Bereich der IT-Beratung gefunden werden.

Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und Scheer Group GmbH

Projektlaufzeit: 01.03.2014 bis 29.02.2016

Automatische Extraktion von Reproduktionsschritten für Software-Fehler aus Benutzer-Interaktionen

Tobias Röhm

ReproFit



Software-Fehler sind ein Ärgernis für Softwarenutzer. Um Fehler zu beheben, müssen Entwickler diese reproduzieren. Die Fehlerreproduktion stellt sicher, dass ein Fehler vorliegt, und hilft die Fehlerursache zu finden. Dazu brauchen Entwickler Informationen über die Reproduktionsschritte eines Fehlers. Allerdings erhalten Entwickler diese Information nur selten, weil sie in Fehlerberichten von Nutzern nicht beschrieben werden und von automatisierten Werkzeugen meist nicht erfasst werden. Deshalb sind Software-Fehler oft schwer oder gar nicht zu reproduzieren und zu beheben.

Um dieses Problem zu adressieren, haben wir den ReproFit-Ansatz entwickelt. ReproFit zeichnet Interaktionen zwischen einem Nutzer und einer Anwendung mit einem Capture/Replay-Werkzeug auf. Dann extrahiert ReproFit Reproduktionsschritte aus aufgezeichneten Interaktionsprotokollen. Dazu werden zwei sich ergänzende Algorithmen verwendet: Delta Debugging und Sequential Pattern Mining. Delta Debugging simuliert automatisiert Teilmengen eines aufgezeichneten Interaktionsprotokolls um so die minimale, fehler-induzierende Teilmenge zu identifizieren. Sequential Pattern Mining ermittelt die gemeinsame Teilmenge von aufgezeichneten Interaktionsprotokollen, die den gleichen Fehler auslösen. Extrahierte Reproduktionsschritte werden von ReproFit als visueller Fehlerbericht verständlich dargestellt. ReproFit wurde in einer Fallstudie erfolgreich evaluiert.

ReproFit hilft Entwicklern, Fehler zu reproduzieren und zu beheben und so die Qualität ihrer Anwendung zu verbessern. Weiterhin beschäftigt sich dieses Projekt mit Privatsphäre-Fragestellungen bei der Protokollierung und Verwendung von Nutzerinteraktionsdaten.

Projektpartner: Technische Universität München und Software AG

Projektlaufzeit: 01.03.2014 bis 28.02.2015

Verarbeitung von inkongruenten Intentionen in systemgestützten, non-kollaborativen Dialogen

Sabine Janzen

SatIN

Erinnern Sie sich an Ihren letzten Dialog? Vielleicht war es ein kollaborativer Dialog mit übereinstimmenden Intentionen der Gesprächspartner, wie z.B. bei der gemeinsamen Lösung eines PC-Problems. Wahrscheinlicher ist, dass Sie Teilnehmer in einem nicht-kollaborativen Dialog waren, in dem nicht übereinstimmende und zum Teil zueinander in Konflikt stehende Intentionen vorhanden waren, z.B. in einem Verkaufsdialog. Nicht-kollaborative Dialoge repräsentieren den Hauptteil der täglichen Kommunikation zwischen Menschen, wurden aber im Rahmen der Dialogplanung bisher nur partiell betrachtet. Ziel des Projektes „SatIN“ war es, die Verarbeitung übereinstimmender und gegensätzlicher Intentionen in Dialogen theoretisch zu erfassen, mit Hilfe empirischer Studien am Beispiel von Verkaufsdialogen zu analysieren, Ergebnisse zu modellieren und durch die Implementation eines prototypischen Dialogsystems zu evaluieren.

Zu den wesentlichen wissenschaftlichen und technischen Ergebnissen des Projekts gehört die Konzeption und formale Beschreibung eines Modells „SatIN“, das auf Basis eines spieltheoretischen Gleichgewichtsansatzes Dialogsysteme befähigt, übereinstimmende und gegensätzliche Intentionen der Gesprächspartner strategisch zu verarbeiten, so dass ein Dialog geplant werden kann, der von allen Teilnehmern als fair empfunden wird.

Die erarbeiteten Forschungsergebnisse wurden in einem webbasierten Retailing Conversational Agent implementiert, der auf Basis verteilter, semantischer Produktinformationen nicht-kollaborative Verkaufsdialoge mit Kunden führt. Abschließend wurde im Rahmen einer empirischen Nutzerstudie mit 120 Probanden insbesondere die wahrgenommene Fairness des Dialogs mit dem Dialogsystem untersucht. Die Ergebnisse zeigen eine positive Beurteilung des Dialogsystems durch die Teilnehmer im Hinblick auf seine Fähigkeit, einen fairen Dialog zu generieren und dabei jeweils auf die Intentionen der Kunden einzugehen, auch wenn gegensätzliche Intentionen bestehen.

Für die wirtschaftliche Praxis ergeben sich durch die Integration eines solchen Dialog-Service neue Gestaltungspotentiale für die Umsetzung von innovativen, webbasierten Verkaufs- und Beratungsservices, z.B. in Online-Shopping Situationen. Im Kontext von Beratungsdialogen im Online-Shopping übernimmt das Dialogsystem die Funktion eines Beraters für den Kunden und stellt zum anderen eine Schnittstelle für die Unternehmenskommunikation dar, d.h. es vertritt während des Dialogs die wirtschaftlichen Interessen des Herstellers oder Händlers. Eine Erweiterung von bestehenden IKT-Systemen um einen solchen Dialog-Service wäre aufgrund der webbasierten Umsetzung des Dialogsystems mit überschaubarem Aufwand möglich.

Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und Scheer Group GmbH

Projektlaufzeit: 01.08.2012 bis 31.07.2014

Collaborative Enrichment of Business Processes

Christina Di Valentin

SCORE

Mitarbeiter verfügen in den seltensten Fällen über gleiche Vorkenntnisse und gleiches Know-How. Auch die Präferenz bei der Nutzung von Software-Applikationen ist von großer Diversität geprägt. Folglich benötigt jeder Mitarbeiter im Rahmen der Bearbeitung von Aufgaben unterschiedliche Informationen und ggf. eine unterschiedliche Unterstützung. Gleichzeitig zeichnet sich in den letzten Jahren vermehrt der Trend des kontextsensitiven Arbeitens („context aware computing“) ab. Wird innerhalb eines Unternehmens bspw. eine neue Software eingeführt, dauert es eine gewisse Zeit, bis dies auf eine Akzeptanz der Mitarbeiter stößt. Zwar werden nach der Einführung der neuen Unternehmenssoftware in den meisten Fällen Schulungen und Fortbildungen angeboten, jedoch wird das in der Schulung erworbene Wissen schnell vergessen, wenn sich Mitarbeiter nicht mit konkreten Problemstellungen befassen. Häufig werden auch in Unternehmen Standards für die Bearbeitung spezifischer Aufgabenstellungen eingesetzt, die meist nicht an den aktuellen Wissensstand von Mitarbeitern angepasst sind. Obwohl Wissen, insbesondere im unternehmensspezifischen Kontext, eine immer wichtigere Rolle einnimmt, werden im Durchschnitt nur etwa 30% des innerhalb eines Unternehmens vorliegenden Wissens tatsächlich genutzt. Folglich ist es sowohl für Mitarbeiter als auch für das Unternehmen wichtig, Wissen bereitzustellen, das an die aktuelle Situation im Arbeitsprozess angepasst ist. Mitarbeiter erzielen hierdurch eine wesentliche Erleichterung und Zeitersparnis im Rahmen ihrer Aufgabenbearbeitung, was für Unternehmen mit einer verringerten Fehlerquote sowie einem erhöhten Wertschöpfungsbeitrag und Kostensenkungen einhergeht.

Ziel von SCORE ist die prototypische Entwicklung einer Softwarelösung, die sowohl unter Berücksichtigung des aktuellen Kontextes von Mitarbeitern als auch unter Berücksichtigung von deren Profildaten den aktuellen Arbeitskontext von Mitarbeitern erkennt und auf auftretende Frage- oder Problemstellungen eingehen kann. Hierfür werden Mitarbeitern Informationen im Sinne von Best Practices angezeigt, die sie im Rahmen der Bearbeitung einer spezifischen Aufgabe innerhalb eines Arbeitsprozessschrittes unterstützen. Die Profildaten, auf die im Rahmen der Empfehlungsfunktionalitäten zurückgegriffen wird, können vom Mitarbeiter selbst in der zu entwickelnden Lösung angegeben werden („explizite Profildaten“). Darüber hinaus werden Profildaten durch Interaktionen, die Mitarbeiter innerhalb der Software durchführen, kontinuierlich angereichert („implizite Profildaten“).

Situationsbedingte Privacy Interfaces

Florian Gall

SiPrI

Die allgegenwärtige Datenerfassung und -Verarbeitung stellt Nutzer und Entwickler vor neue Herausforderungen im Kontext des Schutzes privater und personenbezogener Informationen. Vorhandene Kontrollmechanismen decken die komplexen und individuellen Bedürfnisse von Nutzern nur unzureichend ab, und führen damit zu Kontrollverlust und Misstrauen. Um dem zu begegnen bedarf es Mensch-Maschine-Schnittstellen, die intuitiv verständliche und situationsabhängige Privatsphäre-Kontrollen umsetzen.

Um diese Schnittstellen bereits bei der Entwicklung neuer Dienste zu integrieren ist es wichtig, die Sicht der Softwareentwickler mit einzubeziehen. Im Projekt SiPrI sollen daher die Anforderungen der Nutzer sowie der Softwareentwickler analysiert werden, um darauf aufbauend Entwurfsmuster zu definieren, welche die Entwicklung bedienbarer und effektiver Kontrollmechanismen ermöglichen. Dabei ist das übergeordnete Ziel, den Schutz der Privatsphäre ohne eine übermäßige Einschränkung von innovativen Diensten zu ermöglichen.

Privatsphäre-Einstellungen sind Kontrollmöglichkeiten für Endnutzer, die sich auf die Freigabe von personenbezogenen Informationen für Dritte beziehen. Aktuelle Studien zeigen, dass sich bisherige Mechanismen, insbesondere im Bereich sozialer Netzwerke, als unzureichend herausgestellt haben. Es ist eine Diskrepanz zwischen der Absicht der Nutzer und dem tatsächlichen Verhalten der Dienste zu erkennen, was erhebliches Missbrauchspotenzial durch versehentlich geteilte Daten schafft und in Folge das Vertrauen der Konsumenten in neue Dienste und Produkte mindert. Wird dieser Konflikt nicht aufgelöst, sind negative Auswirkungen auf die Technologieakzeptanz und die Diffusion von Innovationen in Deutschland zu befürchten.

Die Einbeziehung von Kontext-Informationen hebt das Vorhaben von bestehenden Praktiken ab. Diese setzen meist auf statische Einstellungen, welche von den Nutzern im Vorfeld konfiguriert werden müssen, ohne direkten Bezug zur späteren Nutzung. Unsere Forschung soll helfen, Software innovativ zu gestalten ohne dabei die Daten und das Vertrauen ihrer Nutzer zu gefährden. Dies wird über Entwurfsmuster sichergestellt, die empfohlene Kontrollmechanismen bereitstellen. Dadurch können Entwickler den Aufwand der Anforderungsanalyse erheblich reduzieren. Aufbauend auf dem aktuellen Stand der Forschung im Bereich Nutzeranforderungen an Privatsphäre in sozialen Netzwerken werden Entwickler nach Lösungsansätzen und ihrem Verständnis von Privatsphäre befragt. Die daraus abgeleiteten Entwurfsmuster fließen in die Entwicklung eines Prototyps ein, der abschließend in Alltagssituationen evaluiert wird.

Projektpartner: Technische Universität München und Verlagsgruppe Holtzbrinck Publishing Group

Projektlaufzeit: 01.02.2014 bis 28.02.15

Secure Media Cloud

Alexandra Mikityuk

SMeC

Die fortschreitende Verlagerung von Inhalten in die sog. Cloud findet derzeit auch vermehrt bei Mediendiensten, wie Internet-basiertem Fernsehen (IPTV, Fernsehprogramme werden per Internet übertragen), Anwendung. Vor allem die Speicherung von Filmen in – und die Übertragung von Fernsehkanälen aus – zentralen Cloud-Infrastrukturen wird bereits in großem Maßstab eingesetzt und stellt durch die benötigte Bandbreite und neuartigen Inthalteschutzmechanismen neue Anforderungen an die verfügbaren Übertragungskapazitäten des Internets und an die Verschlüsselungstechniken.

In einem nächsten Schritt wird auch vermehrt die Ausführung von Diensten und Applikationen in die Cloud verlagert (siehe Abbildung 1.1).

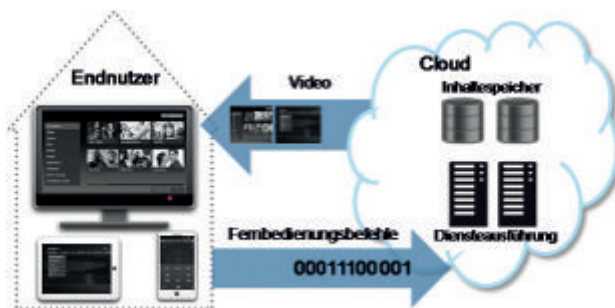


Abbildung 1.1 Verlagerung von Inhaltsspeicher und Dienstauführung in die Cloud

Diese Verlagerung vom lokalen Endgerät (Computer, TV, Smartphone) auf Server in der Cloud erfordert auch hier – wie bereits auf den klassischen Endgeräten – eine geschützte, sichere und standardisierte Ausführungsumgebung. Die serverseitige Ausführung von Diensten wirft eine Vielzahl von bisher ungelösten wissenschaftlich-technisch Fragestellungen auf, wie z.B. bereits angesprochene zur Zeit nicht vorhandene geschützte Umgebungen für die Ausführung der Dienste und neuartige Inthalteschutzmechanismen, die die neuen Nutzungsszenarien in der Cloud berücksichtigen.

Das Forschungsvorhaben „Secure Media Cloud“ setzt genau hier an und untersucht Aspekte der geschützten und sicheren Dienste-Generierung, -Ausführung und -Komposition in der Cloud, sowie des nötigen Inthalteschutzes bei der Übertragung zum Endgerät.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.04.2013 bis 31.10.2014

Semiautomatische Erkennung und Wissensbasis-Integration von Neuigkeiten aus Textdokumenten

Michael Färber

SUITE

Unternehmen stehen zunehmend vor der Herausforderung, die stetig steigende Anzahl an Webdokumenten zu sichten, auszuwerten und schließlich hinsichtlich ihrer Relevanz und Neuartigkeit zu beurteilen. Beispielsweise müssen große Hersteller von Mobiltelefonen aktuelle Trends in der Technologiedomäne entdecken bzw. diese überwachen. Großkonzerne, aber auch mittelständische Unternehmen und Großaktionäre sind auf ein Monitoring der Märkte bzw. auf eine selektive Berichterstattung angewiesen. Die bisher eingesetzten Systeme und Workflows, Neuigkeiten, beispielsweise neue Technologie-Einsatzmöglichkeiten oder Marktveränderungen jeglicher Art, in frisch publizierten Webdokumenten und Nachrichtentexten automatisch zu detektieren und in einem zweiten, oft semi-automatisch vollzogenen Schritt zu bewerten, sind oftmals sehr ineffizient.

Das Projekt SUITE widmet sich daher einem neuen Ansatz zur semantischen Suche nach Neuigkeiten in Textdokumenten. Ziel des Projektes ist die automatische Erkennung von neuen Entitäten sowie von neuen Beziehungen bekannter Entitäten in Textdokumenten. Damit wird es möglich, unmittelbar auf Veränderungen im wettbewerblichen Gefüge und auf Technologie-Neuheiten aufmerksam zu werden und rechtzeitig und adäquat darauf reagieren zu können. Im Rahmen des Projekts kommen Use case-spezifische Wissensbasen zum Einsatz, die die zentrale Rolle der Wissensspeicherung einnehmen. Der Wissensstand der jeweilig eingesetzten Wissensbasis spiegelt den momentanen Wissensstand des Benutzers bzw. der Mitarbeiter in einem Unternehmen wider. Auch andere Ereignisse, wie sie typischerweise in Nachrichtentexten von Nachrichtenagenturen erstmals beschrieben werden, können so modelliert werden. Aufgrund der möglicherweise fachspezifischen Domänen werden spezifische Ontologien gebildet, die auf die Bedürfnisse der Endbenutzer zugeschnitten sind. Die Ergebnisse des Projekts sollen auf Basis von Web-Standards erreicht und auf ihre Flexibilität, Skalierbarkeit und Effizienz in Fallstudien evaluiert werden.

Projektpartner: Karlsruher Institut für Technologie (KIT) und SAP SE

Projektlaufzeit: 01.04.2014 bis 31.03.2016

Usability von Software-Verifikationssystemen: Evaluierung und Verbesserung

Sarah Grebing

USV

Neben den üblichen, in der Industrie verwendeten Methoden, um die Qualität von Software sicherzustellen (insbes. Testen), können formale Methoden, wie etwa Programmverifikation, eingesetzt werden, um die Zuverlässigkeit von Software zu verbessern. Dabei wird mit Hilfe eines Beweissystems geprüft, ob die Implementierung formal spezifizierten Anforderungen genügt. Die Leistungsfähigkeit solcher Beweissysteme hat sich in den letzten Jahren erheblich verbessert. Trotzdem beschränkt sich die Anwendung der formalen Verifikation in der Praxis bisher auf wenige, kritische Bereiche. Dies hat seine Ursache nicht in einer mangelhaften Leistungsfähigkeit der Verifikationssysteme. Vielmehr ist die Usability (Benutzbarkeit) der Verifikationsmethoden und -werkzeuge ein wesentlicher Grund. Usability spielt jedoch in der Entwicklung von Verifikationssystemen bisher nur eine untergeordnete Rolle.

Bei komplexen Beweisaufgaben ist es erforderlich, dass der Nutzer in den Beweisprozess eingreifen muss, um dem System bei der Beweisfindung zu helfen. Im allgemeinen Fall des interaktiven Beweisens wendet das System eine große Anzahl von Beweisregeln auf eine initiale Beweisverpflichtung automatisch an. Wenn das System nicht in der Lage ist den Beweis automatisch zu finden, präsentiert es dem Nutzer den aktuellen Beweiszustand als mathematisch-logische Formel, die mehrere Bildschirmseiten umfassen kann. Der Nutzer muss sich nun in diesem Beweis zurecht finden. Er muss verstehen was passiert ist, an welcher Stelle im Beweisprozess der Beweiser aufgehört hat, was die ihm gezeigte Formel bedeutet und welche weiteren Beweisregeln er anwenden muss. Die Aufgabe einen Beweis interaktiv zu führen erfordert u.a. Strategien zum Problemlösen und Verständnis des Beweises und der Beweisführung. Die Interaktion ist bei den heutigen interaktiven Systemen ein Aspekt des Beweisprozesses, der viel Zeit erfordert und damit die Produktivität des Systems verschlechtert.

Ziel dieses Projektes ist es, die Usability und damit die Akzeptanz von interaktiven Software-Verifikationssystemen zu erhöhen. Voraussetzung dazu ist es, die Usability der Systeme zu evaluieren. Dazu wurden zwei bestehende Verfahren zur Bewertung und Evaluierung von Verifikationssystemen benutzt (Fokusgruppen und Usability Testing) um Methoden und Systeme untereinander zu vergleichen und Mängel aufzudecken. Ausgehend von den Ergebnissen der Evaluierung wurde ein Lösungsansatz für das schnellere Verständnis des automatischen Beweisprozesses entwickelt und mittels Usability Testing evaluiert. Ein Projektergebnis ist eine Evaluierung von Verifikationssystemen in Bezug auf ihre Benutzbarkeit. Ein weiteres Resultat ist eine Identifikation von drei großen Herausforderungen, die ein interaktives Beweissystem bewältigen muss, um für den Menschen benutzbarer zu sein, zusammen mit einer prototypischen Implementierung des Lösungsansatzes.

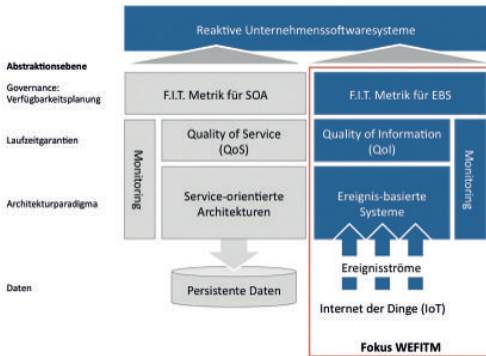
Projektpartner: Karlsruher Institut für Technologie (KIT) und DATEV eG

Projektlaufzeit: 01.03.2013 bis 31.10.2014

Weiterentwicklung der FIT-Metrik für Ereignis-basierte Softwaresysteme

Sebastian Frischbier

WEFITM



Im Internet der Dinge (Internet of Things, IoT) stellt eine Vielzahl unterschiedlichster Datenquellen kontinuierlich Informationen über Realweltereignisse bereit und macht diese für Softwaresysteme nutzbar.

Moderne Unternehmenssoftwaresysteme vereinen daher zunehmend Service-orientierte Architekturen (SOA) und Ereignis-basierte Systeme (EBS) zu hybriden Architekturen, um diese Ereignisströme verarbeiten und dynamisch auf sich ändernde Situationen reagieren zu können. Für den kostenoptimalen Betrieb lose gekoppelter Softwaresysteme müssen auf IT-Managementebene Komponenten und Abhängigkeiten hinsichtlich ihrer Kritikalität für die gesamte Anwendungslandschaft quantifiziert und ggf. angepasst werden. Entsprechende Kritikalitätsmetriken bauen auf Informationen über die von den einzelnen Komponenten zuvor vereinbarten Laufzeitgarantien auf. Laufzeitgarantien werden mittels geeigneter Mechanismen vom System überwacht und durchgesetzt.

Im Gegensatz zu SOA-basierten Systemen existieren für EBS jedoch derzeit keine vergleichbaren Ansätze zur Formulierung, Überwachung und Durchsetzung von Laufzeitgarantien, die relevant für Unternehmenssoftwaresysteme sind.

Ziel des Vorhabens WEFITM ist daher, die Forschungslücke hinsichtlich Laufzeitgarantien für Ereignis-basierte Ansätze in Unternehmenssoftwaresystemen zu schließen und geeignete Formulierungs-, Überwachungs- und Durchsetzungsansätze zu entwickeln. Dies ist notwendig, um entsprechende, für SOA entworfene Kritikalitätsmetriken wie die FIT-Metrik, auch auf EBS anwenden und so EBS in die Governance von Unternehmenssoftwaresystemen integrieren zu können.

Projektpartner: Technische Universität Darmstadt und Software AG

Projektlaufzeit: 01.01.2013 bis 31.03.2015

Partner

Industriepartner:



BOSCH
Technik fürs Leben



SIEMENS

Deutsche Post DHL



Scheer Group
THE INNOVATION NETWORK

 **holtzbrinck**
Publishing Group

Forschungspartner:



TECHNISCHE
UNIVERSITÄT
DARMSTADT

TUM
Technische Universität München

KIT
Karlsruher Institut für Technologie

mpi
max planck institut
informatik



UNIVERSITÄT
DES
SAARLANDES



Deutsches
Forschungszentrum
für Künstliche
Intelligenz GmbH



Fraunhofer
IUK-TECHNOLOGIE

Management:



Förderer:

GEFÖRDERT VOM



Bundesministerium
für Bildung
und Forschung

KONTAKT

Dr. Udo Bub

E-Mail: udo.bub@ictlabs.eu

Tel.: +49 (0)30 34 50 66 90 - 100

Erik Neumann

E-Mail: info@softwarecampus.de

Tel.: +49 (0)30 34 50 66 90 - 150

Maren Lesche

E-Mail: presse@softwarecampus.de

Tel.: +49 (0)30 34 50 66 90 - 140

EIT ICT Labs Germany GmbH
Management-Partner des Software Campus
Ernst-Reuter-Platz 7
10587 Berlin
Geschäftsführer: Dr. Udo Bub