

DEINE IDEE. DEIN IT-PROJEKT. DEINE ZUKUNFT.

Software Campus-Summit 2017

4. Software Campus-Summit

Präsentation der IT-Forschungsprojekte

Der Software Campus bildet die IT-Führungskräfte von morgen aus, indem er Spitzenforschung und Managementpraxis verbindet. Die Teilnehmerinnen und Teilnehmer des Programms sind herausragende Doktorandinnen und Doktoranden sowie Masterstudierende der Informatik und informatiknaher Disziplinen. Sie werden in Führungskräfte trainings weitergebildet und durch erfahrene Mentoren der Industriepartner unterstützt. Außerdem leiten sie ein eigenes Forschungsprojekt, welches vom Bundesministerium für Bildung und Forschung (BMBF) mit bis zu 100.000 Euro gefördert wird.

Bereits zum vierten Mal präsentieren die Teilnehmerinnen und Teilnehmer nun ihre IT-Forschungsprojekte auf einem Software Campus Summit. Aktuelle Teilnehmer, Alumni und Partner des Programms erhalten Einblicke in die Projekte – seien es die Ziele eines gerade gestarteten Vorhabens oder auch Ergebnisse von bereits abgeschlossenen.

Die präsentierten Projekte werden umgesetzt mit den Partnern Technische Universität Berlin, Technische Universität Darmstadt, Karlsruher Institut für Technologie (KIT), Technische Universität München, Universität des Saarlandes, Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI), Fraunhofer-Verbund IuK-Technologie, Max-Planck-Institut für Informatik sowie Robert Bosch GmbH, DATEV eG, Holtzbrinck Publishing Group, Deutsche Post DHL Group, Deutsche Telekom AG, SAP SE, Scheer Holding, Siemens AG, Software AG und werden vom BMBF gefördert.

Inhaltsverzeichnis

- S. 4 **AFAP:** Ausführbare Formalisierungen und Analysen für Prozesse im Web (Tobias Käfer)
- S. 5 **ARaaS:** : Erweiterung einer VaaS-Infrastruktur für Augmented-Reality-as-a-Service (Christian Stein)
- S. 6 **Arcus:** Dezentralisierung der Cloud in Enterprise und ISP Netzwerken (Matthias Rost)
- S. 7 **AWS:** Argumentative Writing Support: Intelligente Unterstützung für argumentatives Schreiben (Christian Stab)
- S. 8 **BodyAnalyzer:** Mobile Erfassung und Analyse von 3D Körpermodellen (Oliver Wasenmüller)
- S. 9 **BONuS:** Bewertung und Optimierung der Nutzerfreundlichkeit von Sicherheitssystemen (Denis Feth)
- S. 10 **CNA:** Connecting the Dots: Auffindung kausaler Zusammenhänge zwischen Nachrichtenartikeln (Nils Reimers)
- S. 11 **DNA:** Die DNA des IoT: Distribute and Aggregate (Niklas Semmler)
- S. 12 **Eko:** Intelligente Entwicklungswerkzeuge zur korrekten Wiederverwendung bestehender Softwarekomponenten (Sven Amann)
- S. 13 **FASTDATA:** Echtzeit-Datenanalyse in Zeiten des Internet of Things (IoT) unter Verwendung von verteilten Hauptspeicher-Datenbanksystemen (Andreas Kipf)
- S. 14 **FeinPhone:** Partizipatorische Feinstaubmessungen mit Smartphones in Szenarien zukünftiger Smart Cities (Matthias Budde)
- S. 15 **FifAKS:** Formale Informationsflußspezifikation und -analyse in Komponentenbasierten Systemen (Simon Greiner)
- S. 16 **GreenFlow:** Entwicklung und prototypische Implementierung des Konzepts ökologischer Prozessmuster am Beispiel IT-orientierter Geschäftsprozesse (Patrick Lübbecke)
- S. 17 **INDIGO:** Konzeption und Entwicklung eines hybriden Ansatzes zur automatisierten Überführung von handgezeichneten Geschäftsprozessdiagrammen in maschineninterpretierbare Formate (Manuel Zapp)
- S. 18 **LinAprom:** Konzeption und Entwicklung eines Systems zur Text-Mining-basierten Unterstützung von Aufgaben des Geschäftsprozessmanagements (Tim Niesen)

- S. 19 **MIRIN:** Computational Main-Memory Databases: Integration von Data-Mining-Prozessen in das Hauptspeicherdatenbanksystem HyPer (Linnea Passing)
- S. 20 **OppEPM:** Opportunistic Exploiting the Power of Mobile Devices –Verfahren zur opportunistischen Nutzung der verfügbaren Potentiale von Mobilgeräten (The An Binh Nguyen)
- S. 21 **PAPMAT:** Ein Privacy Awareness und Privacy Prevention Toolkit für soziale Netzwerke (Frederic Raber)
- S. 22 **PersonalAssistant:** Personalisiertes Assistenzsystem für effiziente Alltags- und Arbeitsunterstützung (Christian Meurisch)
- S. 23 **PersoProfi:** Automatische Vorhersage von Persönlichkeitsmerkmalen fiktiver Charaktere (Lucie Flekova)
- 2.24 **PRODIGY:** Business Process Management using Big Data Driven Predictive Analytics (Nijat Mehdiyev)
- S. 25 **QualS:** Die Qual der intelligenten Standortwahl – Wissensbasierte Entscheidungsfindung über geografische und ökonomische Standortfaktoren (Sebastian Baumbach)
- S.26 **RUMTIME:** Real-Time Usability Improvement auf der Basis von Process Mining (Sharam Dadashnia)
- S. 27 **SADiS:** Smart Attention-Direction Shelf (Denise Kahl)
- S. 28 **SemGo:** Extraktion von semantischen Beziehungen zwischen Geschäftsprozessmodellen zur Unterstützung der Modellintegration im Rahmen der Konsolidierung heterogener betrieblicher Informationssysteme (Philip Hake)
- S. 29 **SONIC:** A new generation of Dynamic Online Social Networks (Sebastian Göndör)
- S. 30 **STEAM:** Data Mining Analytics Framework for Spatiotemporal Data (Bersant Deva)
- S. 31 **SuGraBo:** Suchmaschine für die grafische Benutzeroberfläche (Sven Hertling)

Ausführbare Formalisierungen und Analysen für Prozesse im Web

Tobias Käfer

AFAP

Web-Ressourcen können statische Daten im Web bereitstellen, aber auch dynamische Daten wie Sensorwerte. Das Internet der Dinge und Microservices sind Ansätze neueren Datums, auch Funktionalität in der Form von Web-Ressourcen bereitzustellen. In diesem Kontext wollen wir im Projekt AFAP Programmierung ermöglichen. Semantische Technologien ermöglichen dabei die einfache Integration der Daten verschiedener Quellen. Basierend auf der theoretischen Analyse der relevanten Web-Standards und dem Architekturstil des Web, REST, sowie der empirischen Analyse von Zeitreihen aus dem Linked-Data-Web werden Ansätze zur Beschreibung von Verhalten, z.B. aus der Geschäftsprozessmodellierung, evaluiert, beispielhaft ein Ansatz für die Umsetzung auf der Web-Architektur angepasst und mithilfe von Web-Technologien implementiert.

Projektpartner: Karlsruher Institut für Technologie und SAP SE

Projektlaufzeit: 01.03.16 bis 28.02.2018

Erweiterung einer VaaS-Infrastruktur für Augmented-Reality-as-a-Service

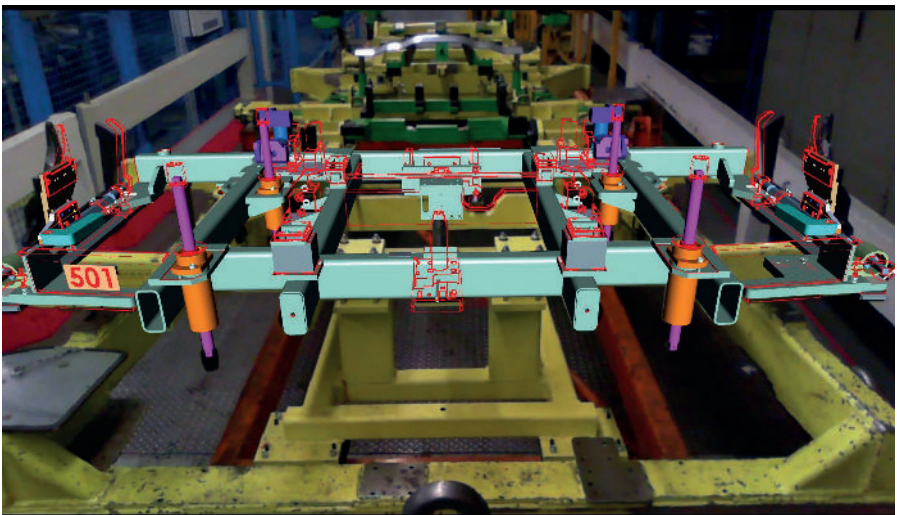
Christian Stein

ARaaS

Im Rahmen der Industrie 4.0 stellen sich für die Computergrafik vor allem zwei Herausforderungen. Zum einen wächst, zum Beispiel im Zuge der fortschreitenden Digitalisierung industrieller Prozesse oder aufgrund der immer höheren Genauigkeit von Scanner-Daten, die Größe der zu visualisierenden Datensätze deutlich schneller als die Leistungsfähigkeit von Prozessen und Grafikkarten. Zum anderen sollen diese Daten heutzutage jederzeit, überall und im besten Falle auf jedem Endgerät visualisierbar sein.

Diese Anforderungen erfüllt die Visualization-as-a-Service-Plattform instant3Dhub des Fraunhofer IGD in dem sie je nach Einsatzszenario die notwendigen Berechnungen automatisch zwischen Server und Client verteilt.

Im Rahmen des Software-Campus-Projektes Augmented-Reality-as-a-Service soll eine ähnliche Funktionalität für AR-Szenarien erreicht werden, indem auch hier verschiedene Berechnungsschritte abstrahiert werden, um diese dann flexibel zwischen Client und Server verschieben zu können. Damit sollen AR-Szenarien auf Basis verschiedener Trackingmethoden auch auf leistungsschwachen Geräten vollautomatisch zur Verfügung stehen.



Projektpartner: Fraunhofer-Verbund IuK-Technologie und Robert Bosch GmbH

Projektlaufzeit: 01.04.2016 bis 31.03.2018

Dezentralisierung der Cloud in Enterprise und ISP Netzwerken

Matthias Rost

Arcus

Cloud Computing, also die Erbringung von Diensten in Rechenzentren, hat in den vergangenen Jahren massiv an Bedeutung gewonnen. Heutzutage werden teils ganze Anwendungen, wie z.B. die Office-Software oder das Enterprise-Resource-Planning System, in die Cloud ausgelagert. Dieser Wandel ist maßgeblich durch die gesteigerte Flexibilität in der Provisionierung begründet: Virtualisierungstechniken erlauben es ohne lange Wartezeit weitere Rechenkapazität in der Form von virtuellen Maschinen zur Verfügung zu stellen.

Die neuen Möglichkeiten in der Provisionierung stellen jedoch auch neue Herausforderungen für die Betreiber von Netzen und Rechenzentren dar. So wächst einerseits die Komplexität von Cloud-Diensten mit der immer verbreiteteren Nutzung derer, was eine automatische Provisionierung der Ressourcen erschwert. Andererseits wächst vor allem im Bereich des Managed Hostings die Erwartungshaltung der Kunden an die erbrachten Dienste in Bezug auf die genaue Erfüllung der vertraglich vereinbarten Dienstgüte.

Vor diesem Hintergrund werden im Arcus Projekt Dienstmodelle entwickelt und betrachtet, welche einerseits detailliert genug sind, um spezifische Kundenwünsche berücksichtigen zu können, und welche andererseits immer noch automatisch von den Cloud-Providern provisioniert werden können. Bezüglich der Kundenwünsche werden im Projekt insbesondere die Skalierbarkeit von Diensten sowie die entsprechende Bepreisung der Ressourcen untersucht. Auf der Seite der Provider werden im Projekt effiziente Algorithmen entwickelt, um entsprechende von Kunden angefragte Ressourcen schnell bereitstellen zu können. Zu diesem Zweck werden im Projekt effiziente Algorithmen mit Gütegarantien für die Einbettung von Diensten entwickelt.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.03.2016 bis 31.08.2017

Argumentative Writing Support: Intelligente Unterstützung für argumentatives Schreiben

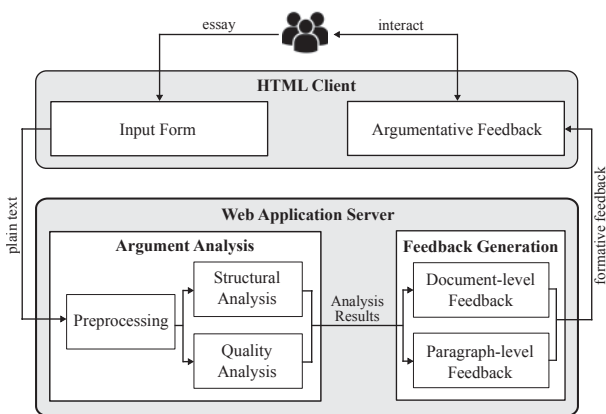
Christian Stab

AWS

Argumentation ist ein wesentlicher Bestandteil der Entscheidungsfindung und vieler Alltagssituationen. Die Konsequenzen von unüberlegten Entscheidungen sind oft weitreichend und können zu fatalen Konsequenzen führen. Daher analysieren Fachleute wie bspw. Juristen, Wissenschaftler oder Entscheidungsträger Pro- und Contra-Argumente systematisch, bevor eine Entscheidung getroffen wird. Allerdings ist das Formulieren von überzeugenden Argumenten eine komplexe Aufgabe und setzt fundierte Argumentationsfähigkeiten voraus.

Um die Argumentationsfähigkeiten von Schülern zu fördern, wurde im Projekt „Argumentative Writing Support“ (AWS) eine neue Technologie zur intelligenten Schreibassistenz entwickelt, die automatisch Feedback zu argumentativen Texten generiert. Die entwickelte argumentative Schreibunterstützung ermöglicht Schülern und Studenten, ihre Argumentationsfähigkeiten zu trainieren und Schwachstellen in ihren Argumenten zu erkennen.

Dazu wurden neue Sprachtechnologien zur automatischen Erkennung und Bewertung von Argumenten in studentischen Aufsätzen entwickelt. Diese ermöglichen die Erkennung von Argumentationsstrukturen, die Identifikation von nicht begründeten Behauptungen oder die Bewertung der Gründe eines Arguments. Somit bieten die im Projekt „Argumentative Writing Support“ entwickelten Methoden nicht nur die Grundlage für intelligente Schreibassistenten, sondern ermöglichen auch neue Anwendungen in Bereichen der Informationssuche oder der Entscheidungsunterstützung.



Projektpartner: Technische Universität München und Holtzbrinck Publishing Group

Projektlaufzeit: 01.03.2015 bis 28.02.2017

Mobile Erfassung und Analyse von 3D Körpermodellen

Oliver Wasenmüller

BodyAnalyzer

Individuelle anthropometrische Maßzahlen sind wichtig für unterschiedlichste Anwendungen (Gesundheitswesen, Ergonomie, Bekleidungsindustrie). Aus 3D Körpermodellen können diese Daten automatisch extrahiert werden. In diesem Projekt wurde ein Körperscanner entwickelt werden, der sich durch seine Präzision, Robustheit und Handhabbarkeit sowie Ergänzung einer Analysekomponente als Body Analyzer von seiner Konkurrenz abhebt. Die benötigte Hardware zur Aufnahme der Daten ist dabei auf eine einzige Kamera reduziert.

Ein visueller Körperscanner rekonstruiert ein geschlossenes 3D Modell einer Person aus unterschiedlichen Kameraansichten. Diese Ansichten können auf der Basis einer Kamera gewonnen werden, indem sich die Person vor dieser Kamera um die eigene Achse dreht, oder, indem die Kamera um die Person herum bewegt wird. Geht man von einer Tiefenkamera aus, so werden aktuelle Filter- und globale, nicht-starre Registrierungsverfahren verwendet, um die aufgenommenen Daten in ein konsistentes Körpermodell zu überführen. Solche Systeme wurden in der Vergangenheit bereits entwickelt und publiziert. Diese Systeme weisen jedoch einige Einschränkungen in Präzision, Detailgrad und Robustheit auf, die den praktischen Einsatz bisher erschweren. In diesem Projekt wurden die Einschränkungen weitestgehend durch gemeinsame Verwendung von Farb- und Tiefeninformationen beseitigt und das System durch Ergänzung einer Analysekomponente zur Extraktion anthropometrischer Maße zu einem kostengünstigen, ohne technische Kenntnisse handhabbaren Body Analyzer erweitert.

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

Projektlaufzeit: 01.03.2015 bis 31.05.2016

Bewertung und Optimierung der Nutzerfreundlichkeit von Sicherheitssystemen

Denis Feth

BONuS

Datensicherheit (Security) ist eine Grundvoraussetzung für den Betrieb und die Nutzung moderner Systeme und stellt seit Beginn der zunehmenden Digitalisierung und Vernetzung ein wichtiges Forschungsgebiet dar. Durch Überwachung und Steuerung der Ressourcennutzung werden Schäden durch Angreifer auf verschiedenen Systemebenen verhindert. Dies hat allerdings auch Schattenseiten. Sicherheitsfunktionen sind häufig schwer zu verstehen, schränken Anwender ein oder führen zu einer zusätzlichen Belastung für die Anwender. Im schlimmsten Fall lehnen Anwender solche Systeme ab oder umgehen sie, womit ein vermeintlich sicheres System unsicher wird. Ziel des Projekts „BONuS“ ist es die Umsetzung von Security und Privacy-Mechanismen aus Anwendersicht zu optimieren. Dazu werden Security und Privacy explizit in den Human-Centered-Design-Prozess integriert. Des Weiteren dient ein Qualitätsmodell der Bewertung der Usability von Sicherheitsfunktionen. Durch das Erreichen dieser Ziele wird ein besseres Verständnis der Beziehung zwischen Security, Privacy und Usability geschaffen. Außerdem werden Entwickler darin unterstützt Sicherheitsmechanismen zielgerichtet auf die Anwender der jeweiligen Domäne abzustimmen um eine höhere Akzeptanz bei den relevanten Nutzergruppen zu erreichen.

Connecting the Dots: Auffindung kausaler Zusammenhänge zwischen Nachrichtenartikeln

Nils Reimers

CNA

Die Informationsüberflutung ist eine zentrale Herausforderung in unserer digitalisierten Welt. Immer mehr Informationen stehen immer schneller zur Verfügung und drohen deren Adressaten in Passivität versinken zu lassen. Um hiermit im privaten als auch im wirtschaftlichen Umfeld umgehen zu können, bedarf es effektiver Methoden zur Strukturierung von Informationen, individuell zugeschnitten auf die Bedürfnisse des Einzelnen.

Das Projekt Connecting News Articles spezialisiert sich dabei auf die Informationsüberflutung im Bereich der Nachrichten. Durch den Wegfall der gedruckten Zeitung sind Online-Nachrichtenportale in der Lage, deutlich öfter und deutlich umfassendere Artikel zu publizieren. Selbst auf nur einem Nachrichtenportal werden teilweise bis zu 100 Artikel pro Woche zu einem einzigen Thema veröffentlicht. Immer auf dem neuesten Stand zu bleiben ist damit fast unmöglich.

Im Rahmen des Projektes Connecting News Articles wurden unterschiedliche Methoden, die eine Unterstützung bei den genannten Herausforderungen anbieten, erforscht und prototypisch umgesetzt. Dazu werden die Informationen aus Nachrichten-Artikeln automatisiert extrahiert und artikelübergreifende Informationen miteinander in Bezug gesetzt. Diese Extraktion und Verlinkung kann verwendet werden, um gezielt die Informationen anzuzeigen, an denen ein Nutzer interessiert ist. Dadurch kann dieser sowohl einen deutlich schnelleren Überblick über Themen erhalten als auch gezielter solche Informationen auffinden, die für ihn von Interesse sind.

Trag's Saddam Hussein, facing U.S. and Arab troops at the Saudi border, today sought peace on another front by promising to withdraw from Iranian territory and release captured soldiers during the Iran-Iraq war. Also today, King Hussein of Jordan arrived in Washington seeking to mediate the Persian Gulf crisis. President Bush on Tuesday said the United States may extend its naval quarantine to Jordan's Red Sea port of Aden to shut off Iraq's last unhindered trade route. In another mediation effort, the Soviet Union said today it had sent an envoy to the Middle East on a series of stops to include Baghdad. Soviet officials also said Soviet women, children and invalids would be allowed to leave Iraq. President Bush today denounced Saddam's "ruthless policies of war," and said the United States is "striking a blow for the principle that might does not make right." In a speech delivered at the Pentagon, Bush seemed to suggest that American forces could be in the Gulf region for some time. "No one should doubt our staying power or determination," he said. The U.S. military buildup in Saudi Arabia continued at fever pace, with Syrian troops now part of a multinational force camped out in the desert to guard the Saudi kingdom from any new threat. In a letter to President Bush, Rafsanjani of Iran, read by a broadcaster over Baghdad radio, said he will begin withdrawing troops from Iranian territory on Friday and release Iranian prisoners of war. Iran said an Iraqi diplomatic delegation was en route to Tehran to deliver Saddam's message, which it said would review "with optimism." Saddam appeared to accept a border demarcation treaty he had rejected in peace talks following the August 1998 cease-fire of the eight-year

Projektpartner: Technische Universität Darmstadt und Holtzbrinck Technology Group

Projektlaufzeit: 01.03.2015 bis 28.02.2017

Die DNA des IoT: Distribute and Aggregate

Niklas Semmler

DNA

Quer durch die deutsche Wirtschaft laufen Anstrengungen um auf einen Industrie 4.0 Ansatz umzustellen. Data-driven Business Models werden evaluiert, Maschinen mit einer Vielzahl von Sensoren ausgestattet und Softwarelösungen zur Verarbeitung von großen Datenmengen entwickelt. Jedoch lassen sich nicht alle Teile der Infrastruktur so einfach skalieren wie Cloud-Ressourcen. Gerade die Backbone-Verbindungen, die mehrere Industriestandorte und Städte verbinden, lassen sich nur langsam aufrüsten, werden dabei jedoch von vielen Nutzern geteilt.

Im Projekt DNA (Distribute and Aggregate) analysieren wir, wie verteilte Datenanalysen flexibel auf verfügbare Bandbreiten reagieren können. Zunächst müssen dafür Metriken erhoben werden, die angeben welche Bandbreiten in der Infrastruktur gegenwärtig vorhanden sind und an welchen Stellen der Datenanalyse Flaschenhalse entstehen. Basierend auf diesen Metriken können verschiedene Optimierungsstrategien evaluiert werden. (1) Die Verteilung der Datenanalyse kann angepasst werden durch eine Migration von bestehenden Prozessen auf andere Standorte. (2) Eine Alternative ist die Aggregation der Daten vor ihrer Übertragung. Die bei dieser Aggregation verringerte Genauigkeit der Daten muss entsprechend mit den Anforderungen der Datenanalyse balanciert werden.

Für das Projekt emulieren wir eine Datenanalyse über mehrere Kontinente. Zu diesem Zweck verwenden wir das an der TU Berlin entwickelte Apache Flink und haben es mit einer Vielzahl von Änderungen an unsere Bedürfnisse angepasst. In dieser Demonstration zeigen wir den Effekt von fluktuierenden Bandbreiten auf die Datenanalyse.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.03.2016 bis 28.02.2018

Intelligente Entwicklungswerkzeuge zur korrekten Wiederverwendung bestehender Softwarekomponenten

Sven Amann

Eko

Wiederverwendung bestehender Softwarekomponenten gehört zum Fundament moderner Softwareentwicklung. Daher müssen Softwareentwickler mit diesen Komponenten und den Best-Practices in deren Verwendung vertraut sein. Eine enorme Herausforderung, bedenkt man die Vielzahl von Frameworks, Bibliotheken und Technologien, die in den letzten Jahren entstanden sind. Die deshalb erforderlichen Weiterbildungen um Mitarbeiter effektiv in Entwicklungsteams einsetzen zu können, sind für viele Firmen sehr kostspielig.

Im Software Campus Projekt KaVE haben mein Kollege Sebastian Proksch und ich Assistenzwerkzeuge für Entwickler entwickelt, die beim Schreiben von neuem Code unter Verwendung bestehender Softwarekomponenten unterstützen. Diese Werkzeuge, oft als „Recommender Systems for Software Engineering“ (RSSE) bezeichnet, helfen Entwicklern Softwarekomponenten kennenzulernen und korrekt einzusetzen. Die Basis hierfür bilden Ansätze des maschinellen Lernens, die wir einsetzen um Verwendungsmuster aus bestehendem Quellcode automatisch zu erlernen. Damit wird das Wissen vieler Entwickler gebündelt und jedem Einzelnen unmittelbar zur Verfügung gestellt.

Heutzutage besteht Softwareentwicklung jedoch nicht nur aus der Entwicklung neuer Produkte, sondern auch in der Pflege bestehender Lösungen. Oft hängen solche Legacy-Systeme von alten Versionen anderer Softwarekomponenten ab, da deren Verwendungsmuster sich über die Zeit weiterentwickelt haben und nach einer Aktualisierung potentiell Fehler auftreten, da der Quelltext weiterhin die alten Verwendungsmuster beinhaltet.

In meinem Software Campus Projekt Eko entwickle ich deshalb ein Assistenzwerkzeug, das Entwickler beim Auffinden und Korrigieren fehlerhafter oder suboptimaler Verwendungsmuster von Softwarekomponenten unterstützt. Mit leichten Änderungen, können die Ansätze des maschinellen Lernens aus dem KaVE-Projekt eingesetzt werden um solche Muster zu identifizieren und Vorschläge zu machen, wie diese verbessert werden können.

Mit meinem Projektpartner DHL IT Services erprobe ich mein Assistenzwerkzeug in einem industriellen Umfeld um herauszufinden, wie gut es Entwickler bei der alltäglichen Arbeit unterstützt. Solche Feldstudien sind im Allgemeinen sehr kostspielig, haben aber einen enormen Wert für die Forschung. Gleichzeitig hilft das Projekt die Arbeitsumgebung von DHL-Entwicklern zu verbessern und ermöglicht ihnen so, qualitativ hochwertige Software noch schneller zu entwickeln und zu warten.

Projektpartner: Technische Universität Darmstadt und Deutsche Post DHL Group

Projektlaufzeit: 01.06.2015 bis 31.05.2017

Echtzeit-Datenanalyse in Zeiten des Internet of Things (IoT) unter Verwendung von verteilten Hauptspeicher-Datenbanksystemen

Andreas Kipf

FASTDATA

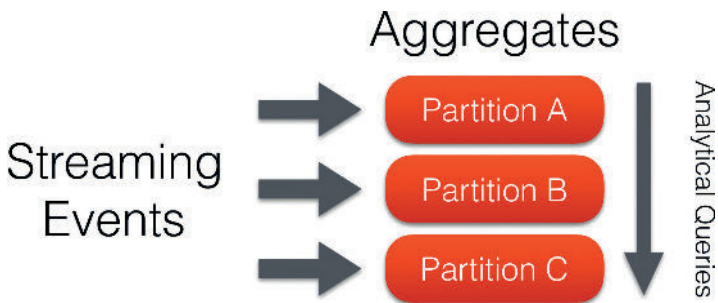
In wenigen Jahren werden mehrere Milliarden Sensoren und Aktuatoren mit dem Internet vernetzt sein. Für viele Unternehmen stellt diese Entwicklung zugleich ein Potential als auch eine Herausforderung dar.

Diejenigen, die das Internet of Things (IoT) in ihren Produkten und Produktionsprozessen zu nutzen wissen, werden dadurch einen enormen Wettbewerbsvorteil erreichen können. Gleichzeitig stehen die Unternehmen jedoch vor der Frage, wie sie die riesigen Datenmengen, die bereits heute durch IoT-Technologien erzeugt werden, verarbeiten sollen. Oft umfassen Analysen sowohl IoT- als auch Geschäftsdaten aus einem relationalen DBMS. Daher ist der Einsatz von traditionellen Streamingansätzen für diese Anwendung ungeeignet, da hierbei die Daten aus verschiedenen Systemen zunächst zusammengeführt werden müssen, was erhebliche Geschwindigkeitseinbußen nach sich zieht. Ziel dieses Projektes ist es, ein verteiltes DBMS zur integrierten Datenverarbeitung zu entwickeln.

Der entwickelte Prototyp basiert auf dem Hauptspeicher-DBMS HyPer, welches an der TU München entwickelt wird. FASTDATA erweitert HyPer um dauerhafte Anfragen, Window-Funktionalitäten sowie eine entsprechende Erweiterung der SQL-Schnittstelle. In einer ersten Studie [1] haben wir HyPer mit modernen Streamingsystemen wie Apache Flink verglichen. Als Beispiel diente eine Anwendung aus der Telekommunikationsbranche, die Analysen auf den sekundenaktuellen Metadaten aller Anrufe durchführt.

Wir haben dabei Unterschiede hinsichtlich der Performanz und der Bedienbarkeit der Systeme festgestellt und Lösungsansätze vorgeschlagen.

[1] Analytics on Fast Data: Main-Memory Database Systems versus Modern Streaming Systems (EDBT 2017)



Projektpartner: Technische Universität München und Siemens AG

Projektlaufzeit: 01.03.2016 bis 28.02.2018

Partizipatorische Feinstaubmessungen mit Smartphones in Szenarien zukünftiger Smart Cities

Matthias Budde

FeinPhone

Von Feinstaub können erhebliche Gesundheitsrisiken ausgehen: Er kann beim Menschen in die Atemwege eindringen, Zellen schädigen und auch toxische Stoffe tief in den Körper bringen.

Die Feinstaubbelastung in Städten wird heute durch teure, statische Messstationen mit schlechter räumlicher und zeitlicher Auflösung überwacht. Um feingranulare Belastungskarten und reaktive Systeme in Smart Cities Szenarien zu ermöglichen, müssten dichte, verteilte Messungen vorgenommen werden. Eine Möglichkeit dafür sind partizipatorische Messungen mit Smartphones. Beim sogenannten „Participatory Sensing“ werden Privatpersonen mit kostengünstigen mobilen Sensoren ausgestattet, etwa integriert in bereits vorhandene Smartphones oder als eigenständige Geräte. Durch die Mobilität der einzelnen Teilnehmer kann eine höhere räumliche Auflösung erreicht werden. Beispiele für die erfolgreiche Umsetzung solcher Ansätze sind etwa Systeme zur Erstellung von Geräuschbelastungskarten oder zur Erfassung von Schlaglöchern, kaputten Ampeln und Verschmutzungen in Städten.

Während solche Projekte meist auf regulären Smartphones und der darin verbauten Sensorik basieren, existieren integrierte Sensoren zur Messung von Feinstäuben in Smartphones noch nicht. Vergangene Arbeiten haben jedoch gezeigt, dass die Hintergrund-Feinstaubbelastung selbst mit äußerst einfachen, bereits relativ kleinen Staubsensoren erfasst werden kann. Prinzipiell ist es auch möglich das Messprinzip dieser Sensoren (Lichtstreuung) an Smartphones mit integrierter Kamera zu adaptieren. Das Projekt FeinPhone hat das Ziel, eine solche neuartige Sensorkomponente für Smartphones zur Messung von Feinstaub zu entwickeln und zu evaluieren.



Projektpartner: Karlsruher Institut für Technologie (KIT) und Siemens AG

Projektlaufzeit: 01.04.2015 bis 31.05.2017

Formale Informationsflußspezifikation und -analyse in Komponentenbasierten Systemen

Simon Greiner

FIfAKS

Anforderungen an moderne IT Systeme enthalten immer komplexere Funktionalitäten kombiniert mit dem Bedarf an hohe Skalierbarkeit der Systeme. Eine verbreitete Technologie, um mit diesen Anforderungen umzugehen, ist eine komponentenbasierte Systementwicklung. Gleichzeitig spielt Sicherheit gegen Angreifer in diesen Systemen eine immer zentralere Rolle, da durch Sicherheitslecks einerseits die Reputation des Systembetreibers beschädigt werden kann, andererseits auch juristische Konsequenzen zu befürchten sind.

Im Rahmen des Mikroprojektes FIfAKS haben wir einen neuen, sprachunabhängigen Sicherheitsbegriff entworfen, der Eigenschaften von komponentenbasierten Systemen ausnutzt und so sehr präzise und ausdrucksstarke Spezifikationen erlaubt. Unser Sicherheitsbegriff basiert auf einer Formalisierung von Informationsfluss, der sogenannten Nichtinterferenz von Information. Hierbei werden Informationen als privat und öffentlich gekennzeichnet und sichergestellt, dass öffentlich zugängliche Ausgaben nicht durch private Eingaben beeinflusst werden. Im Gegensatz zu anderen Nichtinterferenzbegriffen erlaubt unser Ansatz auch das Klassifizieren von Teilinformationen sowie der Existenz von Ereignissen als privat oder öffentlich. Außerdem konnten wir formal nachweisen, dass sichere Komponenten zu sicheren Systemen komponiert werden können.

Aufbauend auf diesem Sicherheitsbegriff haben wir das graphische Spezifikationswerkzeug Palladio so erweitert, dass Sicherheitsspezifikationen in der Entwurfsphase eines Systems festgehalten werden können. Die Spezifikation enthält zunächst auf Domänenebene Informationen, die sich aus dem Einsatzgebiet des Systems ergeben und wird dann auf technischer Ebene verfeinert. Auf technischer Ebene sind Spezifikationen Anforderung an Komponenten, die von den Entwicklern umzusetzen sind. Wir haben unsere Spezifikationssprache benutzt, um für CoCoME (ein Kassen- und Verwaltungssystem für Supermärkte) Sicherheitsspezifikationen bezüglich mehrerer potentieller Angreifer zu formulieren.

Abschließend haben wir ein neues Verfahren zum Testen von Sicherheitseigenschaften in Komponenten entworfen. Das Verfahren baut auf üblichen Fuzz-Test Systemen auf. Hierbei werden zufällig Eingaben für das System generiert, die wir automatisiert so weiterverarbeiten, dass sicherheitsrelevante Fehler offenbart werden. Wir konnten das Verfahren auf eine Implementierung des von uns spezifizierten CoCoME Systems anwenden.

Projektpartner: Karlsruher Institut für Technologie (KIT) und Deutsche Post DHL Group

Projektlaufzeit: 01.03.2015 bis 28.02.2017

Entwicklung und prototypische Implementierung des Konzepts ökologischer Prozessmuster am Beispiel IT-orientierter Geschäftsprozesse

Patrick Lübbecke

GreenFlow

Die Optimierung von Geschäftsprozessen nach ökonomischen Kriterien wie etwa der Durchlaufzeit oder des Ressourceneinsatzes ist eine klassische Problemstellung des Geschäftsprozessmanagements (GPM). In den vergangenen Jahren hat sich der Betrachtungswinkel des GPM neben ökonomischen Faktoren zusätzlich auch noch auf ökologische Faktoren ausgeweitet. Konzepte zur Optimierung von Geschäftsprozessen im Hinblick auf ökologische Aspekte wie etwa des Energieverbrauchs werden unter der Bezeichnung „Grünes Geschäftsprozessmanagement“ zusammengefasst (Grünes GPM). Eine Methode des GPM, die traditionell für die Optimierung ökonomischer Faktoren, mittlerweile jedoch auch immer stärker zur Optimierung ökologischer Faktoren verwendet wird, ist die Geschäftsprozesssimulation. Durch Simulation lässt sich die Auswirkung von Änderungen an einem Geschäftsprozess auf abhängige Variablen wie den Energieverbrauch vor dem Rollout eines optimierten Prozesses untersuchen. Das hat den Vorteil, dass das Risiko des Auftretens von unerwarteten negativen Effekten nach Rollout des revidierten Prozesses minimiert wird. Somit entfiel eine erneute Revision, was gerade bei größeren Unternehmen einen enormen organisatorischen Aufwand darstellt.

Gerade für kleine und mittelständische Unternehmen ist die Anwendung von Simulationsstudien jedoch häufig nicht möglich, da es an der notwendigen Fach- und Methodenkompetenz fehlt. Dennoch sind auch KMUs an der Optimierung ihrer Prozesse in Richtung des Nachhaltigkeitsgedankens interessiert. Die Lösung für dieses Dilemma liegt in der Bereitstellung von Werkzeugen zum Geschäftsprozessmanagement, mit denen KMUs ihre Unternehmensprozesse ohne tiefgreifendes Fachwissen auf Verbesserungspotenzial aus ökologischer Sicht hin untersuchen können.

Das Ziel von GreenFlow ist die Identifikation eines Katalogs von ökologisch nachteiligen Prozessmustern (weakness patterns) für standardisierte Prozessschritte im Verwaltungsbereich. Für diese weakness patterns sollen anschließend ökologisch vorteilhaftere Gestaltungsalternativen vorgeschlagen werden. Diese Muster wiederum werden in der BPM Suite ARIS des Praxispartners Software AG implementiert und ermöglichen KMUs dadurch, eine große Anzahl an Prozessmodellen auf ökologisches Verbesserungspotenzial hin zu untersuchen. Die Erstellung des Katalogs erfolgt durch Analyse bestehender Prozess- und Systemlandschaften mit verschiedenen Verfahren und anschließender Bewertung der ökologischen Auswirkungen einzelner Prozesse und Systeme anhand nachvollziehbarer Kennzahlen.

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

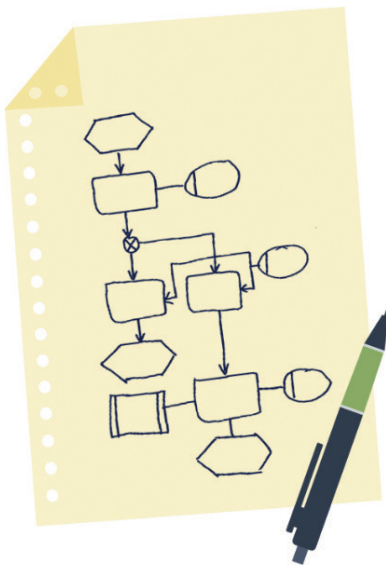
Projektlaufzeit: 01.04.2016 bis 30.09.2017

Konzeption und Entwicklung eines hybriden Ansatzes zur automatisierten Überführung von handgezeichneten Geschäftsprozessdiagrammen in maschineninterpretierbare Formate

Manuel Zapp

INDIGO

Im Geschäftsprozessmanagement haben sich verschiedene Diagrammtypen etabliert, um die Prozesse eines Unternehmens visuell abzubilden. Besonders in der frühen Designphase von Prozessmodellen werden diese Diagramme häufig im Rahmen von Workshops zunächst skizziert. Die Einfachheit des Skizzierens und nicht zuletzt die intuitive Handhabbarkeit tragen zur Beliebtheit dieser Methode bei. Das Skizzieren von Diagrammen hat im Vergleich zum digitalen Modellieren mit dedizierter Software einen großen Nachteil: Skizzen können heute zwar abfotografiert und somit digital abgelegt werden, jedoch nicht in ihrer Semantik erfasst werden, um in vorhandenen digitalen Modellierungstools weiter ausdefiniert zu werden. Im Projekt INDIGO wird genau diese Lücke geschlossen. Auf Basis eines Fotos eines handgezeichneten Prozessdiagramms wird mit Hilfe von Deep Learning ein digitales maschineninterpretierbares Prozessdiagramm erstellt. Diese Diagramme können durch adaptierte Schnittstellen in die Modellierungswerkzeuge des Industriepartners – Software AG – eingeladen, weiterbearbeitet und in die bestehende Prozesslandschaft aufgenommen werden.



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

Projektlaufzeit: 01.04.2016 bis 30.09.2017

Konzeption und Entwicklung eines Systems zur Text-Mining–basierten Unterstützung von Aufgaben des Geschäftsprozessmanagements

Tim Niesen

LinAprom

Die Menge an textuellen Daten in unstrukturierter Form nimmt insbesondere im betrieblichen Umfeld kontinuierlich zu. Sie entstammen einer Vielzahl unterschiedlicher Datenquellen – beispielsweise Emailnachrichten, Ticket-Texten, Social-Media-Kanälen – und unterscheiden sich stark in Bezug auf Semantik, Syntax und Struktur: Während Emailnachrichten üblicherweise syntaktisch vollständige Sätze enthalten, können Daten aus Ticket-Texten inhaltlich weniger konkret und syntaktisch freier gestaltet sein. Gleichzeitig nimmt auch die Menge an Geschäftsprozessmodellen, u. a. bedingt durch die zunehmende Verwendung von Referenzprozessmodellen, beständig zu. Diese Modelle enthalten Texte in Elementbezeichnern (sog. Labels), welche zur Beschreibung von Aktivitäten, eingetretener Zustände und Prozessressourcen dienen. Sie enthalten damit Informationen über die in einem Modell beschriebenen Sachverhalte. Durch die Analyse von textuellen Daten in Dokumenten und Prozessmodellen und die Interpretation der darin enthaltenen Informationen ergeben sich vielfältige Anwendungsfälle für ein ganzheitliches Geschäftsprozessmanagement.

Im Projektkontext werden insbesondere die folgenden beiden Anwendungsszenarien untersucht: Modellierungssupport und Ausführungssupport. Die Unterstützung der Modellierung von Geschäftsprozessen wurde in Form eines innovativen Softwaretools demonstriert, das auf der Grundlage von textuellen Informationen in Prozessmodellen Korrespondenzen zwischen einzelnen Elementen und Teilgraphen aufzeigt und dadurch zu einer effektiven Modellgestaltung beitragen kann. Im Rahmen der Ausführungssupport-Komponente wurde ein handlungsunterstützender Prototyp implementiert, der eingehende IT-Supportanfragen sprachlich analysiert und daraus direkte Vorhersagen zum angestoßenen Prozess ableitet. Hierzu gehören beispielsweise Vorhersagen zur verbleibenden Zeit bis zur Problemlösung und zu den notwendigen Prozessschritten.

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

Projektlaufzeit: 01.04.2015 bis 31.03.2017

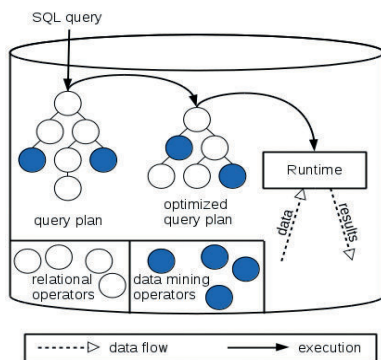
Computational Main-Memory Databases: Integration von Data-Mining-Prozessen in das Hauptspeicherdatenbanksystem HyPer

Linnea Passing

MIRIN

Data Scientists verwalten und analysieren die Datenmengen, die durch Sensoren, Computersysteme und uns Menschen entstehen. Werden für diese Aufgaben Datenbankmanagementsysteme (DBMS) genutzt, lassen sich Anforderungen wie Skalierung, Zugriffsbeschränkung und Datensicherung umsetzen, jedoch unterstützen DBMS bislang kaum komplexe Datenanalysen. Datenanalysensysteme ermöglichen dies, etwa mit Klassifikations- und Vorhersagealgorithmen, erfüllen wegen ihrer simplen Datenhaltungsschicht aber die oben skizzierten Anforderungen nicht. Bei der Nutzung zweier getrennter Systeme erfolgen die Analysen auf veralteten Daten, die Integration der Ergebnisse mit den ursprünglichen Daten ist erschwert, und die IT-Architektur kompliziert. Ziel von MIRIN ist die Integration von Algorithmen zur Datenanalyse in das DBMS, sodass die Vorteile beider Systeme kombiniert werden. So wird die Vision komplexer Echtzeitanalysen auf Geschäftsdaten Realität.

Basis des entwickelten Prototypen ist HyPer, ein hybrides Hauptspeicher-DBMS, welches an der TU München entwickelt wird. HyPer verbindet bereits die Funktionalität eines transaktionalen DBMS mit den Analysefähigkeiten eines Data Warehouse. Mit MIRIN wird ein weiterer Funktionsbereich – komplexe Datenanalysen – hinzugefügt. HyPer nutzt die Möglichkeiten moderner Hardware, etwa Vektorinstruktionen und NUMA-Architekturen, durch Just-In-Time Compilation der Datenbankabfragen effizient aus. Daher lassen sich auch Algorithmen zur Datenanalyse deutlich schneller als in üblichen Datenanalysensystemen ausführen. MIRIN integriert k-Means, Naive Bayes, PageRank und weitere Algorithmen nahtlos in die SQL-Schnittstelle und den Anfrageoptimierer von HyPer.



Projektpartner: Technische Universität München und SAP SE

Projektlaufzeit: 01.03.2016 bis 28.02.2018

Opportunistic Exploiting the Power of Mobile Devices –Verfahren zur opportunistischen Nutzung der verfügbaren Potentiale von Mobilgeräten

The An Binh Nguyen

OppEPM

Die Potenziale hinsichtlich des Einsatzes mobiler Geräte sind vielfältig. Zu den modernen Mobilgeräten gehören z.B. Smartphones, Tablets, Laptops, Raspberry Pi etc. Ein Mobilgerät kann nicht nur Informationen speichern und verarbeiten; sondern durch eine Reihe von Sensoren kann ein Mobilgerät auch verschiedene Informationen von der Umgebung erfassen. Bedingt durch die äußerst dynamische Entwicklung insbesondere des Smartphone-Marktes, sind diese Geräte heute allgegenwärtig. Allerdings werden deren Ressourcen und Leistungsfähigkeiten häufig nicht hinreichend ausgenutzt. Das Identifizieren und Nutzen von solchen ungenutzten Ressourcen stellen ein großes Potential dar.

Das Ziel des OppEPM-Projekts ist es, diese Ressourcen-Nutzung zu erhöhen. Genauer gesagt soll eine Plattform entwickelt werden, mit deren Hilfe freie Ressourcen identifiziert und zum Durchführen von Aufgaben genutzt werden können. Die durchzuführenden Aufgaben sollen in einem sogenannten Workflow den Geräten zugewiesen werden, welche am besten zur Ausführung der Aufgabe geeignet sind. Für das Gesamtsystem sollen keine zusätzlichen Berechnungsknoten eingebracht werden, das System soll stattdessen bereits vorhandene Ressourcen integrieren („Opportunistic Ressource Usage“).

Um dieses Ziel zu erreichen, sollen verschiedene Eigenschaften der Mobilgeräte betrachtet werden, da die Eigenschaften eines Mobilgeräts die Service Discovery sowie die Ausführung der zugewiesenen Aufgaben beeinflussen können. Beispiele für solche Eigenschaften sind die Kommunikationseigenschaften (z.B. verfügbare Netze, Bandbreite), die zur Verfügung stehenden Ressourcen, die Mobilität sowie die Geräteklassen. Dabei ergibt sich die Herausforderung, dass nicht nur die Eigenschaften eines einzelnen Gerätes, sondern der aggregierte Kontext des an einem Workflow partizipierenden Geräts betrachtet werden muss. Eine weitere Herausforderung ist die hochdynamische Umgebung, die ständige Änderung der Netzwerktopologie erzeugt. Solche Kontextänderungen führen dazu, dass Anpassungen hinsichtlich der Service Discovery und der Aufgabenzuweisung zur Laufzeit vorgenommen werden müssen. Im Projekt werden Verfahren und Algorithmen erarbeitet, welche diese Herausforderungen effizient lösen.

Projektpartner: Technische Universität Darmstadt und Siemens AG

Projektlaufzeit: 01.03.2016 bis 31.12.2017

Ein Privacy Awareness und Privacy Prevention Toolkit für soziale Netzwerke

Frederic Raber

PAPMAT

Immer mehr Firmen setzen soziale Netzwerke zur internen Kommunikation ein. Oftmals gelangen ungewollt sensible Informationen an andere Netzwerkmitglieder, oder sind sogar für alle Öffentlich zugänglich. Ziel des Projektes war die Entwicklung einer Software, die zum einen den Nutzer über die momentane Privatsphäre-situation aufklärt (Privacy Awareness), aber zum anderen auch ungewollte Datenverbreitung verhindert, in dem die Privatsphäre Einstellungen automatisch angepasst werden (Privacy Management).

Die Basis des „Awareness“-Teils besteht aus einer Visualisierung, welche Datenflüsse zwischen den Benutzern als Datenflussdiagramm darstellt. Benutzer werden automatisiert zu Gruppen zusammengefasst, welche die Knoten des Graphen darstellen. Die Kanten zwischen den Knoten repräsentieren einen möglichen Datenfluss zwischen den beiden Knoten, wobei die Dicke der Kante die Wahrscheinlichkeit der Datenweitergabe widerspiegelt. Es ist möglich in die einzelnen Gruppen (etwa „Firmenkollegen“) immer tiefer hinein zu „zoomen“, die dann wiederum in mehrere Untergruppen aufgeteilt werden (z.B. „Marketing“, „Vertrieb“).

Der Teil des „Privacy Management“ nutzt verschiedene Datenquellen um die Privatsphäre Einstellungen semi-automatisch einzustellen. Jede Einstellung ist auf den Benutzer zugeschnitten, dessen Persönlichkeit anhand eines speziell entwickelten Fragebogens erfasst wird. Basierend auf Studiendaten errechnet ein machine-learning gestütztes System anhand der Persönlichkeitswerte und des Themas des Posts für jede Freundesgruppe des Benutzers einen sog. „Privacy Level“. Anders als Facebook, welches nur zwei Privacy Level (zeigen/ verbergen) anbietet, kennt PAPMAT fünf Level. Somit kann z.B. nur der Text eines Posts ohne Bildmaterial oder ohne Kommentare angezeigt werden. Studien zeigten dass die Genauigkeit der Vorhersage im Mittel max. um einen Level abweicht, und die Vorschläge von den Benutzern als sinnvoll, z.T. sogar besser als ihre eigenen Einstellungen, betrachtet werden.

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

Projektlaufzeit: 01.03.2015 bis 29.02.2016

Personalisiertes Assistenzsystem für effiziente Alltags- und Arbeitsunterstützung

Christian Meurisch

PersonalAssistant

Die Nutzung von mobilen Endgeräten stieg in den letzten Jahren rasant an und ist aus unserem Alltag nicht mehr wegzudenken. Basierend darauf zeichnet sich ein Trend im Bereich digitaler Assistenzsysteme ab. Assistenzsysteme sind digitale Helfer, die Nutzern im Alltag unterstützend zur Seite stehen. Google Assistant (ehem.: Google Now) und Apple's Siri sind Beispiele, die u.a. helfen den Parkplatz wieder zu finden oder an Termine unter Berücksichtigung der Verkehrslage erinnern, aber nicht auf eine effiziente Erreichung von Nutzerzielen oder einer proaktiven Alltagsunterstützung des Nutzers zugeschnitten sind.

Im Rahmen des Projekts „PersonalAssistant“ soll ein auf den Nutzer personalisiertes und proaktives Assistenzsystem für die Unterstützung im Alltag als auch bei der Arbeit entwickelt werden. Als erster Anwendungsfall werden Studenten als Zielgruppe unterstützt. Jedoch sollen sich die entwickelten Konzepte und das modulare Grundsystem auf andere Zielgruppen, wie beispielsweise auf Büroarbeiter zur Arbeitsunterstützung oder auf pflegebedürftige Menschen zur Zustandserkennung als auch zur Alltagsunterstützung übertragen lassen.

In diesem Vorhaben wird der Schwerpunkt auf die Entwicklung von unaufdringlichen Assistenzkonzepten für die proaktive Unterstützung von Nutzern basierend auf deren Ziele gelegt. Dazu werden Fragestellungen u.a. bezüglich der Erwartungen eines Nutzers an ein solches System als auch bezüglich des geeigneten Zeitpunkts zur Bereitstellung der Assistenz auf Grundlage von Nutzermodellen adressiert. Im Vergleich zu bereits existierenden – meist reaktiven – Systemen, wie Google Assistant oder Apple's Siri, wird dadurch ein Mehrwert geschaffen und das Potenzial derartiger Systeme sowohl weiter aufgezeigt als auch aktiv mitgestaltet.

Projektpartner: Technische Universität Darmstadt und SAP SE

Projektlaufzeit: 01.03.2016 bis 28.02.2018

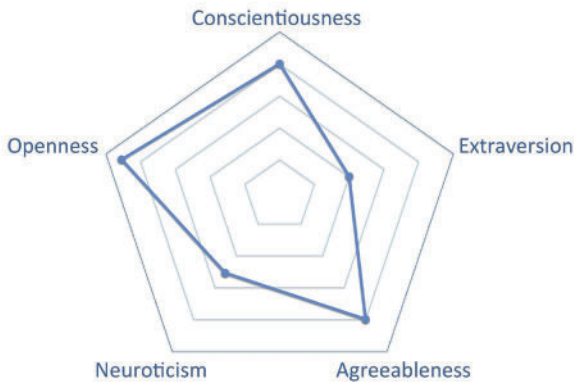
Automatische Vorhersage von Persönlichkeitsmerkmalen fiktiver Charaktere

Lucie Flekova

PersoProfi

Der Anwendungsschwerpunkt unseres Projektes liegt in der Vorhersage der Persönlichkeitsmerkmale der Figuren: Ist Harry Potter mutig? Ist Hermine klug? Haben sie viele Freunde oder sind sie eher einsame Charaktere? Sind sie offen für neue Erfahrungen oder wirken sie eher introvertiert? Unsere Hypothese ist, dass die Popularität eines Buches, zu großen Teilen, davon abhängt, wie sehr der Leser mit den Hauptfiguren sympathisiert. Wenn er sich mit der Geschichte des Helden identifizieren kann, könnte er auch andere Bücher mit ähnlichen Hauptfiguren bevorzugen. Basierend auf den semantischen Informationen, die wir aus den Büchern gewinnen können, wie beispielsweise die Persönlichkeitsmerkmale der Protagonisten bzw. deren Emotionen, wird es möglich, sehr spezifische Empfehlungen zu generieren.

Das Forschungsziel des Projekts liegt daher in der Entwicklung von Methoden zur sprachtechnologischen Verarbeitung von Buchinhalten, die es den Maschinen ermöglichen, Geschichten in einer solchen Tiefe zu interpretieren, dass der Computer in der Lage ist, Folgerungen hinsichtlich der Persönlichkeit der Personen zu ziehen, so wie es ein menschlicher Leser tun würde. Um das zu erreichen, bauen wir nicht nur auf Forschungsteilbereiche der maschinellen Sprachverarbeitung (wie Emotionserkennung oder die Kennzeichnung von semantischen Rollen), sondern auch auf Grundlagen der Forschung in der traditionellen Psychologie, um genau zu sein auf dem Gebiet der Persönlichkeitsprofilierung.



Projektpartner: Technische Universität Darmstadt und Holtzbrinck Publishing Group

Projektlaufzeit: 01.03.2015 bis 28.02.2017

Business Process Management using Big Data Driven Predictive Analytics

Nijat Mehdiyev

PRODIGY

The ability to monitor the industrial processes proactively is one of the main differentiators for firms to stay competitive. It is of a great importance to ensure that the production activities will run in a desired manner by avoiding deviations from designed structure in runtime. Since the firms adopt similar products and identical technologies, high-performance business processes are one of the last points of differentiation. In the PRODIGY research project we explore the possibilities of integrating the Business Process Management, Predictive Analytics and Complex Event Processing approaches to leverage the big data for making more precise data driven decisions.

Complex Event Processing is a suitable technology to enable the real-time fully-automated decision making; however, this event based system lacks the modeling capability and requires the event patterns to be provided by domain experts. Nevertheless, vast amount of data in the presence of intelligent monitoring systems may complicate even make it infeasible for domain experts to identify relationships among diverse data sources and underlying processes. In the PRODIGY research project, the bilateral integration of Predictive Analytics and Complex Event Processing is examined in different settings such as application of rule-based machine learning algorithms to identify the complex event patterns as IF-THEN rules or directly integrating the machine learning algorithms in CEP engines.

Furthermore, embedding predictive modelling and business analytics to the daily business processes in an enterprise is a crucial issue to create the value. Therefore, the PRODIGY research project also investigates the integration of both Predictive Analytics methods and Complex Event Processing technologies with the Business Process Intelligence, through the data generated by Process Aware Information Systems. Such a combination should advance the ability of the enterprises to gain new insights and reduce the gaps between decisions, insights and actions. The project also explores the combination possibilities of human intelligence with machine learning for a complicated decision making in the context of business processes where the automation is not possible and high domain expertise is required. The final purpose of the PRODIGY is the incorporation of the knowledge obtained from business processes through predictive analytics to optimize decisions about the planning activities in the manufacturing environment.

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

Projektablaufzeit: 01.04.2016 bis 30.09.2017

Die Qual der intelligenten Standortwahl – Wissensbasierte Entscheidungsfindung über geografische und ökonomische Standortfaktoren

Sebastian Baumbach

QuIS

Jedes Unternehmen steht mindestens einmal vor der Wahl eines geeigneten Standorts, welcher entscheidenden Einfluss auf den betriebswirtschaftlichen Erfolg besitzt. Dabei gilt es zahlreiche geografische, ökonomische, ökologische und sozioökonomische Faktoren des Standorts zu berücksichtigen und mit den Anforderungen des Unternehmens in Beziehung zu setzen.

Durch die Globalisierung und Digitalisierung vergrößern sich jedoch nicht nur das geografische Suchfenster, sondern auch die verfügbaren Datenmengen. Deren Größe und Komplexität wächst darüber hinaus exponentiell, so dass es selbst für Experten zunehmend immer schwieriger wird, alle relevanten Daten für die Standortwahl zu überblicken. Aus diesem Grund werden in der Praxis sukzessiv Standorte selektiert und andere Regionen damit ausgeschlossen, wobei deren Auswahl nach subjektiven Kriterien erfolgt. Die Standortentscheidung wird demnach nicht konsequent auf Basis aller zugrundeliegenden Informationen getroffen werden.

Im Projekt QuIS wurden Data-Mining-Ansätze der KI eruiert und auf diesen Anwendungsfall übertragen. Das Ziel dieses Projekts war der Entwurf und die prototypische Umsetzung eines Verfahrens, mit dem eine wissensbasierte und objektive Entscheidungsfindung über die Standortfaktoren erfolgen kann. Die zur Verfügung stehenden Datenquellen wurden, unter Berücksichtigung der Standortfaktoren, nach geeigneten Standorten untersucht. Das Verfahren kann dabei einerseits effizienter mit großen Datenmengen umgehen und andererseits Zusammenhänge in den Daten identifizieren, welche in bisherigen Analysen verborgen geblieben sind. Mittels wissensbasierter Entscheidungsfindung können in Zukunft Unternehmen bei der strategischen Standortwahl unterstützt und Big Data in diesem Kontext angewandt werden.

Projektpartner: Deutsches Forschungszentrum für Künstliche Intelligenz (DFKI) und DATEV eG

Projektlaufzeit: 01.06.2015 bis 31.05.2016

Real-Time Usability Improvement auf der Basis von Process Mining

Sharam Dadashnia

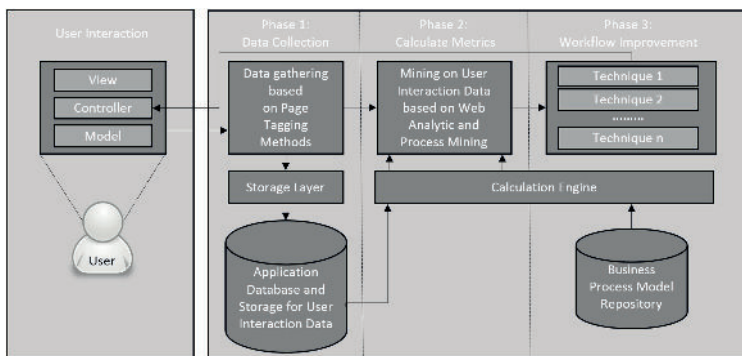
RUMTIME

Die Gebrauchstauglichkeit bzw. Usability von betrieblicher Anwendungssoftware spielt eine entscheidende Rolle bei der Auswahl entsprechender Systeme und stellt ein wichtiges Differenzierungsmerkmal dar. Die Usability solcher Softwaresysteme zu testen und fortlaufend zu verbessern ist oft mit hohen Kosten und Aufwänden verbunden. Um Einsparungspotenziale zu entwickeln, wurden im Rahmen des Projektes RUMTIME Methoden und Konzepte zur dynamischen Verbesserung der Usability von Software während der Laufzeit sowie Analysen zur langfristigen Verbesserung der Usability auf Basis von Process-Mining-Methoden entwickelt. Zentrale Ergebnisse dieses Projektes waren unter anderem die Entwicklung eines Real-Time Usability Improvement Frameworks zur Ad-hoc-Auswertung von Nutzerinteraktionsdaten sowie die darauf aufbauende dynamische Verbesserung auf die individuellen Bedürfnisse der einzelnen Nutzer. Das entwickelte Konzept wurde in drei einzelne Phasen unterteilt:

Phase 1: Die Nutzerinteraktionsdaten werden mittels eines Tracking-Skriptes zu einer In-Memory-Datenbank weitergeleitet. Unter anderem werden Informationen über die Systemnutzung des Nutzers gesammelt, welche die Basis für die späteren Analysen bilden.

Phase 2: Mithilfe der gesammelten Daten können entsprechende Metriken berechnet und zur Usability-Verbesserung der Softwareanwendung herangezogen werden. Hierbei werden größtenteils prozessbezogene Metriken mittels Process Mining berechnet.

Phase 3: Diese Phase bietet eine dynamische Anpassung bzw. Verbesserung der Anwendung auf der Basis der zuvor berechneten Metriken und Ergebnisse aus den Process-Mining-Analysen.



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

Projektlaufzeit: 01.04.2015 bis 31.03.2017

Smart Attention-Direction Shelf

Denise Kahl

SADiS

Im Rahmen des Software Campus-Projektes „Smart Attention-Directing Shelf“ wurde ein Framework entwickelt, welches es ermöglicht, die visuelle Aufmerksamkeit am Einkaufsregal – angepasst an den Kunden und dessen Bedürfnisse – auf relevante Informationen zu lenken. Durch die gezielte Ansteuerung von Aktuatoren am Einkaufsregal ist es hierdurch möglich, Kunden auf interessante Produkte hinzuweisen. Neben dem Kontext und dem Benutzerprofil fließen auch die Reaktionen des Kunden mit in das System ein.

Zu Demonstrationszwecken wurde ein Einkaufsregal, welches über 16 Produktfächer verfügt, aufgebaut und technisch instrumentiert. Mit Hilfe von an den Regalfächern angebrachten Aktuatoren werden verschiedene Reize ausgelöst, welche die visuelle Aufmerksamkeit des Betrachters auf die entsprechenden Produkte lenken, wie z.B. das Beleuchten des zugehörigen Produktfaches. Als Aktuatoren dienen hierbei Smartphones, deren Blitzlichter zum Ausleuchten der Fächer und deren Bildschirme als digitale Preisschilder genutzt werden.

In Benutzerstudien wurden sowohl offensichtliche Reize getestet, welche vom Betrachter erkannt werden, wie z.B. das Blinken des gesamten Preisschildes, als auch subtile Reize, die nur unbewusst wahrgenommen werden, wie z.B. leichte Helligkeitsschwankungen auf einem kleinen Teilbereich des Preisschildes.

Zur Bestimmung des richtigen Zeitpunktes der Reizaussendung sowie der Reaktion des Kunden auf den Reiz wird die Blickrichtung des Kunden zugrunde gelegt. Für die Ermittlung der Blickrichtung werden sowohl ein mobiler Eyetracker als auch ein Motion Capturing System eingesetzt. Die Auswertung der Blickrichtung erfolgt in Echtzeit, sodass die Reaktion des Kunden zur nächsten Reizermittlung direkt wieder in das System einfließen kann.

Umgesetzt wurden in diesem Projekt hierdurch die Anpassung der Reizintensität abhängig von der Entfernung des aktuellen Fixationspunktes des Kunden sowie – je nach Kundenreaktion auf den ausgesandten Reiz – die Anpassung der Reizart.

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

Projektlaufzeit: 01.03.2014 bis 31.08.2015

Extraktion von semantischen Beziehungen zwischen Geschäftsprozessmodellen zur Unterstützung der Modellintegration im Rahmen der Konsolidierung heterogener betrieblicher Informationssysteme

Philip Hake

SemGo

Geschäftsprozesse bestimmen heutzutage maßgeblich die Aufbau- und Ablauforganisation von Unternehmen und stellen einen nicht zu verkennenden Wettbewerbsfaktor dar. Durch das breite Aufgabenspektrum des Geschäftsprozessmanagements, welches Prozesse evolutionär über einen Lebenszyklus betrachtet, entstehen innerhalb von Unternehmen vermehrt Varianten und Versionen von Prozessmodellen. Grund hierfür ist die kontinuierliche Optimierung von Prozessen und den zugehörigen Prozessmodellen. Durch die erhebliche Anzahl an Prozessmodellen in Unternehmen besteht bei der Entwicklung neuer Prozesse analog zu Referenzprozessmodellen ein umfangreiches Potenzial zur Wiederverwendung. Ziel des Projektvorhabens SemGo ist es daher eine effiziente Wiederverwendung von Prozessen und der mit den Prozessen verknüpften Softwaremodulen umzusetzen. Grundlage für eine erfolgreiche Wiederverwendung und ist ein umfangreiches Repositorium an Prozessmodellen. SemGo ermöglicht eine automatisierte Analyse und einen automatisierten Vergleich von Geschäftsprozessmodellen innerhalb des Repositoriums. Die Analysen von SemGo unterstützen primär Prozessdesigner bei der Prozesserstellung durch sprachsemantische Suchen und Vergleiche in bestehenden Prozessmodellen. SemGo greift dabei auf die in Prozessmodellen hinterlegten sprachlichen sowie prozesssemantischen Informationen zurück. Zur Verarbeitung der Informationen werden Techniken des Natural Language Processing (NLP) adaptiert. Um Varianten oder abweichende Prozessmodell-Bausteine innerhalb des Repositoriums zu identifizieren werden zudem strukturelle Hierarchisierungen und Zerlegung von Prozessmodellen vorgenommen.

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

Projektlaufzeit: 01.04.2016 bis 30.09.2017

Sebastian Göndör

SONIC

Im Rahmen des Projekts Sonic wurde die Problematik von Lock-in-Effekten in aktuellen Online Social Network (OSN) Services adressiert, welche durch proprietäre Datenformate und Schnittstellen entstehen. Die Folgen dieses Effekts sind, dass Nutzer an einen Anbieter gebunden werden, so dass sich im Fall eines gewünschten Anbieterwechsels soziale Profildaten wie Foto-galerien oder Texte nicht automatisiert zu anderen Anbietern migrieren lassen. Darüber hinaus unterbinden die proprietären Protokolle und Datenformate eine freie Kommunikation zwischen Plattformen verschiedener Anbieter. Die Folge ist eine Vielzahl von zentralisierten Insellösungen, welche getrennt voneinander existieren.

Um diese Probleme proprietärer OSN-Architekturen zu adressieren, wurde im Rahmen des Projekts der gleichnamige Kommunikationsstandard Sonic mitsamt Protokollen und Datenformaten konzipiert, welcher eine heterogene Föderation von OSN-Plattformen ermöglicht. Die REST-basierten Protokolle basieren hierbei auf existierenden Mikrostandards und gewährleisten hiermit eine größtmögliche Kompatibilität zu bestehenden Plattformen. Sonic unterstützt hierbei ein sogenanntes Core Featureset, welches die von der Mehrzahl existierender OSN-Plattformen unterstützten Funktionalitäten abdeckt. Darüber hinaus bietet Sonic die Option, weitere Funktionalitäten über eine Extension-API zu implementieren. Um Nutzern die Möglichkeit zu bieten, den Anbieter oder die Plattform zu wechseln, wurde eine Migrationsunterstützung integriert. Hierbei wird basierend auf den von Sonic definierten Datenformaten ein automatisierter Wechsel der genutzten OSN Plattform ermöglicht, ohne dass hierbei Daten oder Verbindungen zu anderen Nutzern verloren gehen. Um eine möglichst effiziente Umsetzung in bestehenden OSN Plattformimplementierungen zu gewährleisten, wurde ein quelloffenes Software-SDK implementiert. Hierdurch wird es neuen und bestehenden Plattformen ermöglicht, einer offenen und heterogenen OSN-Föderation beizutreten.



Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.03.2014 bis 31.07.2016

Data Mining Analytics Framework for Spatiotemporal Data

Bersant Deva

STEAM

In den letzten Jahren hat sich der Trend zur Nutzung und Analyse von großen Datenmengen immer stärker ausgeweitet und gewinnt in immer mehr Unternehmen eine größer werdende Bedeutung für wichtige Entscheidungsprozesse. Die Analyse von komplexen, großen und heterogenen Datensätzen erfordert neue Herangehensweisen zur Gewinnung von Erkenntnissen, welche von bisherigen Business-Intelligence Werkzeugen bisher nicht geleistet werden können.

Gleichzeitig, verfügen Unternehmen über immer mehr kontext-bezogene Daten, welche für das Angebot von Dienstleistungen und Produkten eine größere Relevanz erlangen. Einen Teil dieser kontext-bezogenen Daten machen räumlich-zeitliche Daten aus, welche sich durch ihre Heterogenität, ihr Volumen und ihre Komplexität von anderen Datentypen absetzen. Dabei sind die Quellen für räumlich-zeitliche Daten vielfältig, von Bluetooth-, WLAN- und Mobilfunknetzen, hin zu dienstorientierten mobilen Applikationen, die den Standort von Mobilgeräten nutzen. Die Bedeutung von räumlich-zeitlichen Daten nimmt weiter zu mit sich stetig bewegendem und vernetzten Fahrzeugen, wie Autos, Bussen und Lieferwagen, und neuerdings auch Drohnen, fahrenden Robotern und anderen autonomen Vehikeln.

Das Ziel von STEAM ist die prototypische Konzeption und Entwicklung einer Plattform zur skalierbaren verteilten Analyse von großen räumlich-zeitlichen Datenmengen. Die aus den Analysen gewonnenen Erkenntnisse werden dabei so aufbereitet, dass Sie neben der Visualisierung direkt als Informationseingang für weiterführende Dienste genutzt werden können. STEAM ermöglicht die Beantwortung von Fragen in einer Reihe von Anwendungsfällen mit Bezug auf Ort und Zeit. Zum Beispiel, können mit Hilfe des Ansatzes von STEAM, Fragen zur zielgesteuerten orts- und zeitbasierten Kommunikation mit Kunden, zu möglichen Engpässen in der Logistik oder zur Gebäude- und Stadtplanung beantwortet werden.

Projektpartner: Technische Universität Berlin und Deutsche Telekom AG

Projektlaufzeit: 01.04.2015 bis 30.09.2017

Suchmaschine für die grafische Benutzeroberfläche

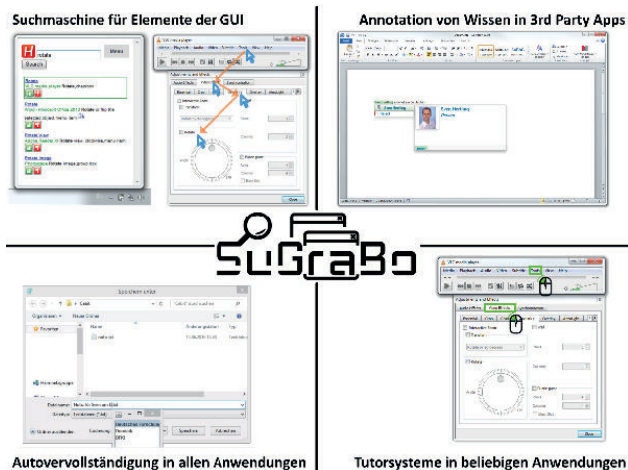
Sven Hertling

SuGraBo

Wie kann ich ein Bild in Gimp oder Photoshop verbessern? Welche Software kann Videos rotieren? Was sollte man bei einer Fehlermeldung 0xc004f050 machen? Die heutigen komplexen Softwaresysteme können oftmals nur schwer von Benutzern über eine grafische Benutzeroberfläche bedient werden.

Im Projekt SuGraBo (Suchmaschine für die grafische Benutzeroberfläche) soll ein Framework entwickelt werden, um Modelle von beliebiger (insb. existierender) Software zu erstellen. Basierend darauf werden verschiedene Assistenzen angeboten, die das Erlernen und den Umgang mit der Software fördern. Die Hauptanwendung ist eine semantische Suche auf den Elementen der Benutzeroberfläche (Schaltflächen, Menüs etc.). Wenn ein Ergebnis ausgewählt wurde, können sogenannte In Application Tutorials angeboten werden. Darüber hinaus können Textassistenzsysteme über Anwendungsgrenzen hinweg gepflegt und genutzt werden. Neben kontextabhängiger Autovervollständigung, können beispielsweise Hintergrundinformationen in Texten eingeblendet werden.

Innerhalb des 4. Summits des Software Campus sollen erste Prototypen und insbesondere die Textassistenzsysteme vorgestellt werden.



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

Projektlaufzeit: 01.03.2016 bis 31.08.2017

Partner

Industriepartner:



BOSCH



Deutsche Post DHL
Group



holtzbrinck
Publishing Group

Scheer
HOLDING

SIEMENS

software AG

Forschungspartner:



Deutsches
Forschungszentrum
für Künstliche
Intelligenz GmbH



Fraunhofer
IUK-TECHNOLOGIE



Karlsruher Institut für Technologie



max planck institut
informatik



TECHNISCHE
UNIVERSITÄT
DARMSTADT

Technische
Universität
München



UNIVERSITÄT
DES
SAARLANDES

Management:

Förderer:



GEFÖRDERT VOM



Bundesministerium
für Bildung
und Forschung

KONTAKT

Dr. Udo Bub

E-Mail: udo.bub@softwarecampus.de

Tel.: +49 30 34506690 - 100

Erik Neumann

E-Mail: erik.neumann@softwarecampus.de

Tel.: +49 30 34506690 - 124

Susanne Kegler

E-Mail: presse@softwarecampus.de

Tel.: +49 30 34506690 - 129

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

