

LDVZ – Nachrichten

Herausgeber:
Landesamt
für Datenverarbeitung und Statistik
Nordrhein-Westfalen

Redaktion:
Bianca Oswald,
Alfons Koegel

Kontakt:
Landesamt für Datenverarbeitung
und Statistik NRW
Postfach 10 11 05
40002 Düsseldorf,
Mauerstraße 51
40476 Düsseldorf

Telefon:
0211 9449-01
Telefax:
0211 442006
Internet:
<http://www.lds.nrw.de>
E-Mail:
poststelle@lds.nrw.de

Auflage:
1 200

© Landesamt
für Datenverarbeitung
und Statistik NRW,
Düsseldorf, 2005
Vervielfältigung und Verbreitung,
auch auszugsweise, mit Quellen-
angabe gestattet.

Bestell-Nr. Z 09 1 2005 52

ISSN 1616-377X

. . . in Kürze

**Unix Server
werden „erwachsen“** **2**
Dr. Frank Dillmann

**Integriertes
Parlamentsinformationssystem
Parlamentsspiegel
Ein Gemeinschaftsprojekt
der deutschen Länderparlamente** **3**
Guido Köhler (Landtag NRW),
Brigitte Corves

Schwerpunktt Themen

**Wahlen in NRW
Die Unterstützung der Landes-
wahlleiterin durch das LDS NRW** **5**
Dr. Thomas Pricking

**Landtagswahl 2005
Interaktive Kartendarstellung
der Wahlergebnisse** **10**
Anja Neumann,
Georg Stahl, Christoph Rath

**Aktuelle Trends
in der Spracherkennung** **13**
Dieter Bittner

**JUDICA/TSJ
Neue Zusammenarbeit
mit dem Justizministerium** **17**
Kirsten Rölleke, Claudia Schröder

**Aufwandsschätzungen
in Softwareprojekten** **20**
Dr. Jan Mütter, Ulrich von Hagen

**Risikomanagement
in Softwareprojekten** **28**
Ulrich Andree, Dr. Jan Mütter

**Projektmanagement
in der Praxis
oder: Wo steht mein
Projekt?** **35**
Thomas Reiff (IBM),
Dr. Jan Mütter

**Server Based Computing
Auf dem Weg in die Zukunft** **44**
Sascha Bittau

Index 2000 – 2005 **53**

Unix Server werden „erwachsen“

Unix Server sind im LDS NRW seit mehr als 10 Jahren im Einsatz. Ein Grund, diese Systeme zu installieren, war, sich der für damalige Verhältnisse „neuen“ Welt zu öffnen. In den vergangenen Jahren ist der Zuwachs an Verfahren, die Systemtechnik im Unix-Umfeld erfordern, nicht nur der Anzahl nach erheblich gewachsen, es werden mittlerweile auch eine Reihe „unternehmenskritischer“ Anwendungen damit betrieben. Somit steigt auch die Anforderung nach erhöhter Verfügbarkeit dieser Systemtechnik.

Die Hardwarehersteller reagieren auf diesen Trend, indem sie im proprietären Unix-Bereich Hardwarelösungen anbieten, die man bisher nur aus dem Großrechnerumfeld kannte. Gekoppelt mit Softwarelösungen wie modernen Clustertechnologien können mittlerweile Verfügbarkeiten von 99,999 % angeboten werden. Dies bedeutet einen statistischen Ausfall von rund 5 Minuten in einem Jahr.

Solche Verfügbarkeiten haben ihren Preis. Als Landesbetrieb muss das LDS NRW kostendeckend und Kosten sparend arbeiten. Somit ergibt sich die Aufgabe unter Kosten-Nutzen-Betrachtungen die für das Verfahren notwendigen Verfügbarkeiten in einer bezahlbaren Form zu realisieren. Die Bereitstellung von Hardware ist dabei allerdings nur ein Aspekt.

Für zentrale Unix-Datenbanken wurde im Jahr 2002 ein erster Hochverfügbarkeitsserver unter dem Betriebssystem Solaris in Betrieb genommen. Es handelt sich um eine Primepower 1000 der Firma Fujitsu-Siemens, die sich sehr bewährt hat und in den vergangenen Jahren keinen nennenswerten, hardwarebedingten Ausfall zu verzeichnen hatte.



Unix-Hochverfügbarkeitsserver

In diesem Jahr wird nun als zentrales Internetsystem für Unix-Datenbanken eine Sunfire 4800 der Firma Sun ihren Betrieb aufnehmen. Beide Systeme wurden nach EU-weiten Ausschreibungen beschafft. Diese Server haben eine vom Hersteller garantierte Verfügbarkeit von Hardware und Betriebssystem von mindestens 99,9 %. Dies bedeutet ein Ausfall von weniger als 9 Stunden in einem Jahr, was nach unserer Einschätzung für Anforderungen der Landesverwaltung an Verfahren im LDS NRW ausreichend ist. Es wird kein automatisches Cluster eingesetzt, da dies die Kosten (vor allem durch den Betreuungsaufwand) erheblich vergrößern würde. Allerdings verfügen diese Systeme über eine so genannte Back-Up Domäne. Dies kann man als „Cold Stand By“ auffassen und gewährleistet beim Ausfall von Teilen des Produktsystems ein Wiederanlaufen nach wenigen Minuten bis maximal einer Stunde. Hierzu sind manuelle Eingriffe notwendig.

Die skizzierten Systeme verfügen über Funktionalitäten, welche früher dem

Großrechnerbereich vorbehalten waren. Augenscheinlich wird das auch in der Tatsache, dass die Firma Fujitsu-Siemens seit wenigen Jahren Großrechner mit dem Betriebssystem BS2000 anbietet, die auf einer Solaris-Systemarchitektur und Unix-Hardware aufbauen.

Die Funktionalitäten dieser Hochverfügbarkeitsserver sind im Einzelnen:

- Partitionierung des Systems (Betrieb einzelner, absolut voneinander getrennter Domänen)
- dynamische Rekonfiguration (Hinzuschaltung von Ressourcen zu einzelnen Partitionen im Betrieb)
- selbstheilende Prozesse
- komplette Redundanzen aller Bauteile (nur das Chassis ist einfach ausgelegt)

Das LDS NRW bietet zur Abdeckung unterschiedlicher Kundenwünsche eine maßgeschneiderte Palette von Servern und Systemen in unterschiedlichen Verfügbarkeitsklassen für bedarfsgerechte wirtschaftliche Lösungen an. Neben den oben skizzierten High-End Unix Lösungen verfügt das LDS NRW auch über ein umfangreiches LINUX-Angebot, welches in den letzten Jahren vor allem im Bereich des webhosting zu einem der Hauptschwerpunkte der Systemtechnik geworden ist.

Die dabei verwendete Serverarchitektur zeichnet sich durch einfache und kostengünstige x86-Systeme aus, mit denen im RZ-Betrieb des LDS NRW gute Verfügbarkeiten erreicht werden. Hierbei sind neben dem einfachen Betrieb auch bedarfsgerechte Sonderlösungen wie Cluster, Failover aber auch virtuelle Systeme möglich.



Dr. Frank Dillmann
Tel.: 0211 9449-2680
E-Mail: frank.dillmann@lds.nrw.de

Integriertes Parlamentsinformationssystem Parlamentsspiegel

Ein Gemeinschaftsprojekt
der deutschen Landesparlamente

Der Parlamentsspiegel ist ein Informationssystem für die Politik in den deutschen Landesparlamenten, der alle parlamentarischen Initiativen, d. h. Drucksachen und Plenarprotokolle, auswertet. Seit 1957 gibt es den Parlamentsspiegel, anfangs als vierzehntägige Lieferung von Karteikarten. Er weist von 1980 bis 1997 nur die überregional bedeutsamen parlamentarischen Initiativen (z. B. Gesetzentwürfe, Anträge, Anfragen, Entschlüsse, Plenarprotokolle und Drucksachen) als Quellen nach und bietet hierzu die Dokumente (soweit öffentlich uneingeschränkt zugänglich) im Originaltext komplett an. Beginnend mit dem Jahr 1997 werden im Parlamentsspiegel alle Initiativen ohne Einschränkungen sukzessive dokumentiert. Der nachfolgenden Aufzählung ist zu entnehmen, ab wann die Daten für die einzelnen Landesparlamente nachgewiesen werden:

- Bayern (ab 1997)
- Niedersachsen (ab Beginn der 14. Wahlperiode März 1998)
- Sachsen-Anhalt (ab Beginn der 3. Wahlperiode Mai 1998)
- Mecklenburg-Vorpommern (ab Beginn der 3. Wahlperiode Oktober 1998)
- Nordrhein-Westfalen (ab November 1998)
- Brandenburg (ab Beginn der 3. Wahlperiode September 1999)
- Sachsen (ab Beginn der 3. Wahlperiode Oktober 1999)
- Thüringen (ab Beginn der 3. Wahlperiode Oktober 1999)
- Berlin (ab Beginn der 14. Wahlperiode November 1999)
- Schleswig-Holstein (ab Beginn der 15. Wahlperiode März 2000)
- Hamburg (ab Beginn der 17. Wahlperiode Oktober 2001)

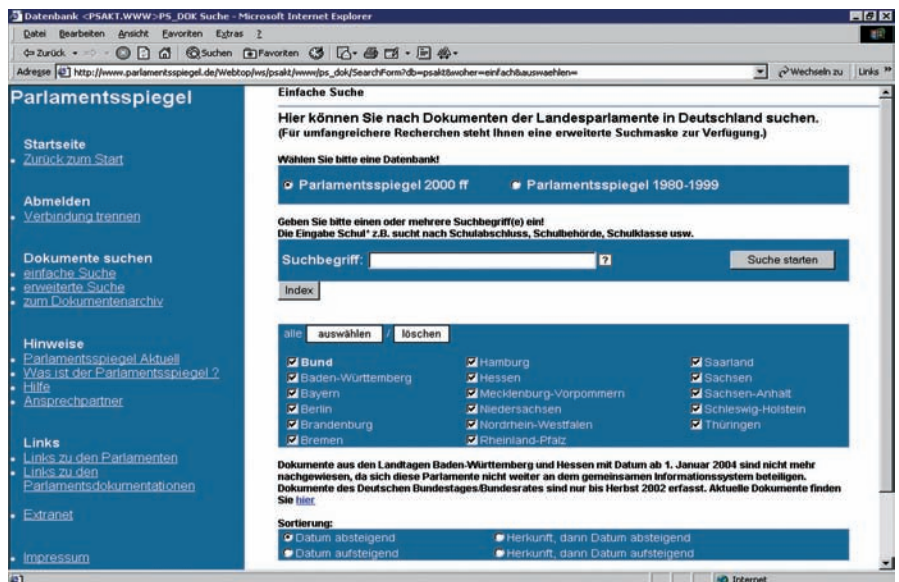


Abb. 1: Einfache Suche

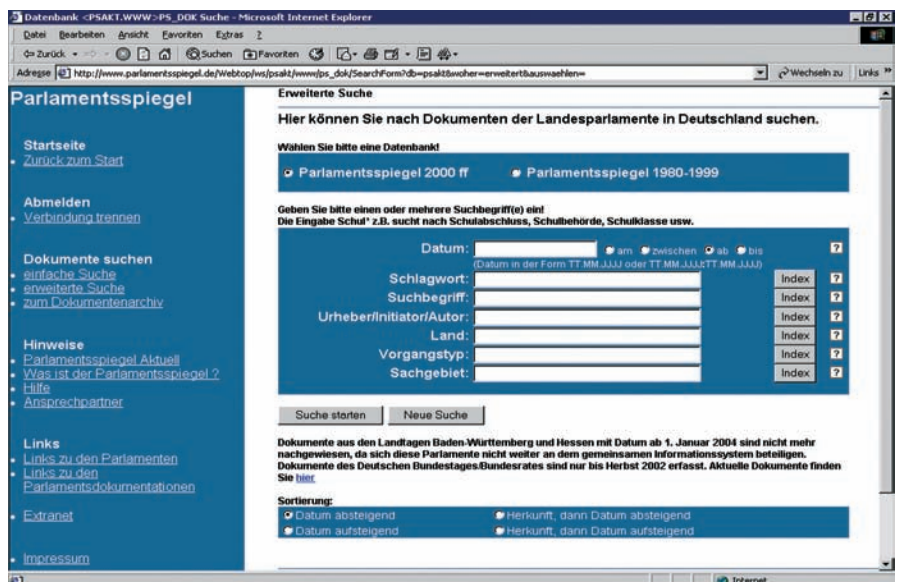


Abb. 2: Erweiterte Suche

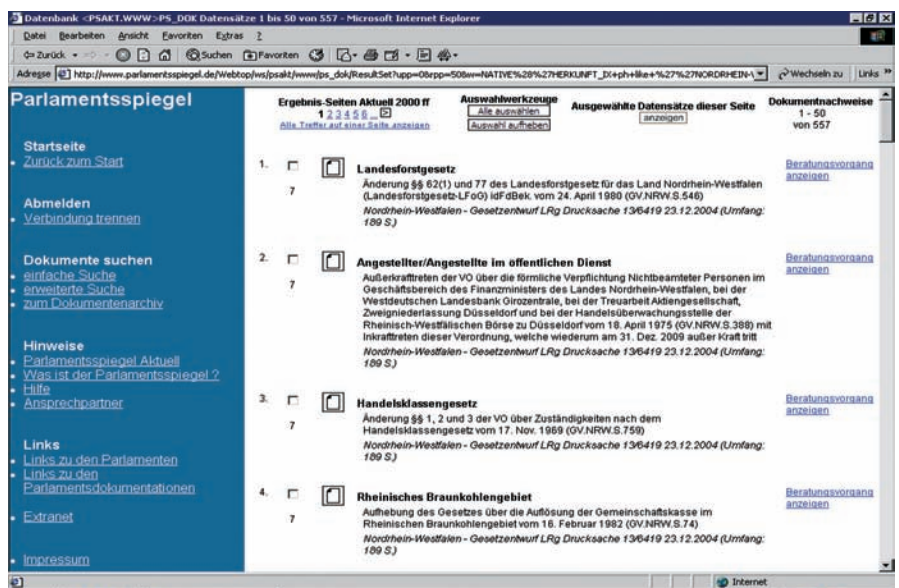


Abb. 3: Ergebnisanzeige

LDVZ-Nachrichten

... in Kürze

- Rheinland-Pfalz
(ab Beginn der 14. Wahlperiode
Mai 2001)
- Bremen
(ab Beginn der 16. Wahlperiode
September 2003)
- Saarland
(ab Beginn der 13. Wahlperiode
September 2004)
- Die Landtage von Baden-Württemberg und Hessen beteiligen sich seit 2004 nicht weiter am Parlamentspiegel. Ihre Dokumente sind nur für den Zeitraum von 1980 bis 2003 vorhanden.

Seit 1980 ist das Informationssystem automatisiert. Die inhaltlichen Auswertungen erfolgten beim Landtag Nordrhein-Westfalen. Ab 1997 wurde das integrierte Informationssystem Parlamentspiegel weiter entwickelt, d. h. die Daten der einzelnen Landtage wur-

den nach und nach direkt als XML-Datei als so genanntes Austauschformat, an den Landtag NRW geliefert, so dass doppelte Auswertungen entfielen. Zusätzlich ist der Zugriff auf die Dokumente teils per Link auf die Internetangebote der einzelnen Landtage direkt möglich, teils liefern die Landtage die Dokumente zur Einspeicherung in das Dokumentationssystem des Landtags Nordrhein-Westfalen.

Als Retrievalsystem wird Livelink Collection Server® der Firma opentext (vormals BASISplus®) und zur Erstellung der Internetoberflächen WEBtop, JAVA und JAVA Server Pages eingesetzt. Für die Suche sind zwei verschiedene Suchmasken als so genannte einfache und erweiterte Suche programmiert worden. Als Ergänzung für

die Datenverwaltung und für sehr spezifische Aufgaben wird eine interne Expertensuchmaske angeboten, die nur den Informations- und Dokumentationsstellen der Landesparlamente zur Verfügung steht. Diese Oberflächen sind barrierefrei gestaltet.

Der Parlamentsspiegel wird regelmäßig aktualisiert und ist unter der Adresse www.parlamentsspiegel.de abrufbar.



*Guido Köhler
Landtag NRW
Tel.: 0211 884-2443
E-Mail: guido.koehler@landtag.nrw.de*



*Brigitte Corves
Tel.: 0211 9449-6305
E-Mail: brigitte.corves@lds.nrw.de*



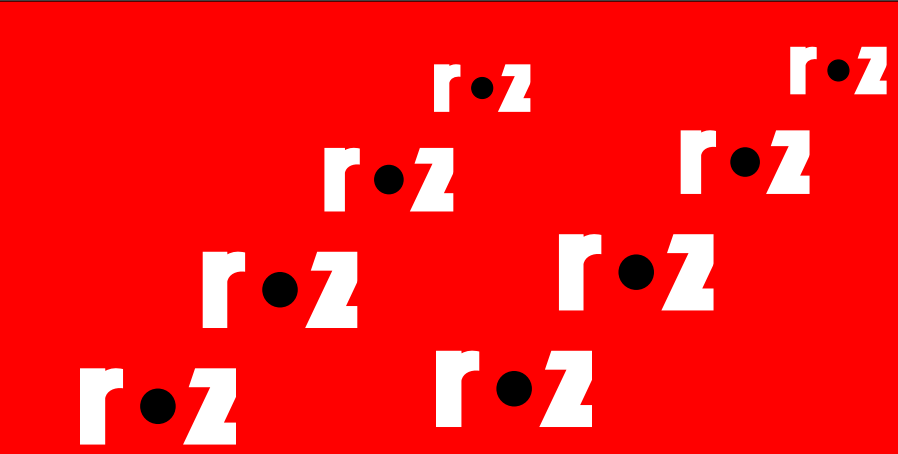
Unser Angebot

Sie sind interessiert, Ihre Daten auf zentralen Systemen zu speichern?

Unsere Stärken sind Ihre Vorteile:

- sicher
- performant
- skalierbar
- wirtschaftlich

Wir freuen uns auf Ihre Aufgabenstellung und unterbreiten Ihnen gern unser Angebot.



Kontakt

Speicherbereitstellung
Telefon: 0211 9449-2180
E-Mail: rechenzentrum@lds.nrw.de

Weitere Informationen finden Sie unter:
<http://lv.landesintranet.lds.nrw.de/datenverarbeitung/rechenzentrum/speicherlp/index.html>

Speicherbereitstellung im Rechenzentrum. Eine Dienstleistung des Landesamtes für Datenverarbeitung und Statistik Nordrhein-Westfalen. 2005





PRE3.001.2005-08

© LDS NRW, Düsseldorf, 2005

Wahlen in NRW

Die Unterstützung der Landeswahlleiterin durch das LDS NRW

„Jede Stimme zählt!“ Mit diesem Slogan bemühen sich die Parteien seit langem, die Wählerinnen und Wähler zum Urnengang zu mobilisieren. Ändert man den Slogan nur geringfügig ab in „Jede Stimme wird gezählt!“, so wird kurz und bündig umschrieben, was seit Jahrzehnten nach der Schließung der Wahllokale um 18 Uhr in den Gemeinden und Kreisen des Landes und im Innenministerium in großer Eile und ebensolcher Präzision erfolgt: Die Ermittlung des vorläufigen amtlichen Endergebnisses der jeweiligen Wahl. Der nachfolgende Beitrag skizziert am Beispiel der Landtagswahl 2005, welche Vorbereitungen und Umsetzungen, welche Probleme und Lösungen für die erfolgreiche Bereitstellung der Ergebnisse erforderlich sind und umgesetzt werden. Anders als bei Bundestags- oder Europawahlen müssen die Ergebnisse ja nicht nur an den Bundeswahlleiter weitergeleitet werden, der dann die Sitzverteilung und Listenplätze übermittelt. Bei der Landtagswahl ist die Errechnung der Sitzverteilung, der Überhang- und Ausgleichmandate zusätzlich die genuine Aufgabe des LDS NRW.

Die Vorbereitung

Die Vorbereitungen für die Landtagswahlen am 22. Mai 2005 setzen sich aus zwei Phasen zusammen. In der ersten, langfristig angelegten Phase hat das Wahlteam ein Basissystem zur Durchführung der Wahlen (WaBaSys) entwickelt, mit dem ein großer Schritt in Richtung Vereinheitlichung und Wiederverwertbarkeit der für die Ergebnisermittlung erforderlichen Software-Module getan worden ist. Arbeitsprozesse, die wahlübergreifend strukturell identisch sind, können mit nur geringen Anpassungen wiederholt eingesetzt werden.

- Die Erfassung der Wahlkreisergebnisse: Bei jeder Wahl sind unterschiedliche Daten in das System aufzunehmen. Während die komplexeste Struktur bei den Kommunalwahlen mit den Ergebnissen für die Vertretungen der kreisfreien Städte und den Kreistagen, für die Oberbürgermeister, Bürgermeister und Landräte vorliegt, müssen bei den Bundestagswahlen Erst- und Zweitstimme gezählt werden. Die Europawahl und die Landtagswahl sind mit der Erfassung nur einer Stimme recht einfach. Bei jeder Wahl treten unterschiedliche und andere Parteien an, die auch nicht flächendeckend kandidieren, sondern teilweise nur in einzelnen Wahlkreisen mit Bewerbern auftreten, oder auch keine Landeslisten einreichen. Trotz dieser

Unterschiede ermöglicht WaBaSys mit Hilfe von hinterlegten Datenbanktabellen eine weitgehend flexible Steuerung der jeweils erforderlichen Eingabemasken und Plausibilitätsprüfungen.

- Der Druck diverser Wahlveröffentlichungen: Seit alters her ist es Brauch, dass vor der Wahl eine Veröffentlichung mit den Ergebnissen der vorhergehenden Wahlen, den auf die jeweils aktuellen Wahlkreise umgerechneten Ergebnissen anderer Wahlen sowie einer Auswahl soziodemografischer Merkmale erscheint, die zur Information des Wahlvolkes und der Politik dient und historische Vergleiche und politisch-strategische Planungen ermöglicht. Am Morgen nach dem Wahlabend muss ein Heft (Auflagenhöhe 1 300 Exemplare) mit den vorläufigen amtlichen Endergebnissen, mit Listen der gewählten Bewerber und der über Listen in das Parlament gezogenen Kandidaten und Rankings der Parteiergebnisse bereitgestellt werden; dieses Heft wird nach dem Vorliegen des endgültigen Ergebnisses neu aufgelegt. Außerdem werden die Ergebnisse nicht nur mit dem regionalen Bezug der Wahlkreise, sondern in einem weiteren Heft auch auf Ebene der administrativen Gebietseinheiten (kreisfreie Städte und kreisangehörige Gemeinden oder Kreise) dargeboten. Das Bewerberverzeichnis vor der Wahl und die Ergebnisse der repräsentativen Wahlstatistik – Ergebnisse nach Alter und Geschlecht der Wählerinnen und Wähler – runden mit zwei weiteren Heften die Wahlpublikationen ab. WaBaSys unterstützt das Wahlteam nachhaltig bei der Druckerstellung und -steuerung, ermöglicht es doch die Parametrisierung aller Formate, Layouts, Tabellenstrukturen, Seitenangaben, Inhaltsverzeichnisse und dergleichen mit wenigen Mausklicks und bindet die Vorgaben in automatische Druckjobs ein. Da sich diese Vorgaben im Zuge der Tests sehr gut erproben lassen, braucht am Wahlabend der Druck nur noch angestoßen zu werden, so dass unmittelbar verwendbare Druckvorlagen oder beim Einsatz entsprechender Filter PDF-Dokumente entstehen.
- Die Bereitstellung der Ergebnisse im Internet: Seit den ausgehenden 1990er-Jahren werden die Wahlkreisergebnisse unmittelbar nach der fachlich-formalen Prüfung und Freigabe auf einen Webserver überspielt und stehen dort im Internet der interessierten Öffentlichkeit zum Abruf bereit. Das Angebot ist bisher bewusst schlicht gehalten gewesen und hat sich auf die Darstellung der Daten in einfacher Tabellenform beschränkt, um die Nutzer, die teilweise nur über geringe Zugangsbreiten ins Internet verfügen (analoges Modem oder ISDN-Anschluss), nicht mit

unnötigen Warte- und Ladezeiten zu belasten. Mit der zunehmenden Verbreitung von leistungsfähigen Netzen (z. B. DSL) scheint der Zeitpunkt gekommen zu sein, die Darstellung der Ergebnisse im Web zu verbessern. So wurden für die Landtagswahl 2005 erstmals grafische Präsentationen in Form von Diagrammen und dynamischen Kartendarstellungen in das Web-Angebot eingebunden¹⁾.

Aufgrund der Standardisierung mit WaBaSys konnten die Freiräume geschaffen werden, die für die Umsetzung der jeweils aktuellen und neuen Anforderungen der zweiten Vorbereitungsphase erforderlich sind, ohne dass es – wie in der Vergangenheit häufig geschehen – zu erheblichen Belastungen im Wahlteam kommt. Diese Phase ist bestimmt von mehreren Teilprozessen:

- Ermittlung der Vergleichszahlen: Sofern von Wahl zu Wahl keine Änderungen im regionalen Zuschnitt der Wahlkreise vorgenommen worden sind, können die „alten“ Wahldaten sofort für Vergleichszwecke herangezogen werden. In der jüngeren Vergangenheit war das kaum mehr der Fall: Im Zuge der angestrebten Verkleinerungen der Parlamente oder der Umsetzung der gesetzlichen Regelungen zu den Wahlkreisgrößen werden die Wahlkreise neu geschnitten, so dass mit Unterstützung der Kommunen und Kreise die Vergleichsdaten auf Ebene der Stimmbezirke zusammengetragen und neu aggregiert werden müssen. Die Umrechnung von Wahldaten unterschiedlicher Wahlen ist zumeist ebenfalls nur auf der Basis der Stimmbezirke möglich und bedarf auch der Mitwirkung der Kommunen.
- Aufbau der kartographischen Grundlagen: Werden die Wahlkreise verändert, müssen auch die georeferenziel-

len Daten angepasst werden, indem vorhandene Karten digitalisiert oder vorhandene Bezugswerte umgesetzt werden.

- Aufbau der Datenbank mit Steuerungsinformationen: Damit alle Programme auf einer korrekten Zuordnung der Bewerber und Parteien zu den Wahlkreisen oder Listenplätzen, der Vergleichsergebnisse zu den aktuellen Wahldaten sowie der geographischen Basisdaten zu den Ergebnissen aufbauen, wird die Wahldatenbank kontinuierlich mit den hierfür erforderlichen Regelungs- und Plausibilisierungsvorgaben befüllt.

Mit dem Umzug des Wahlteams aus dem LDS NRW in das Innenministerium, dem Sitz der Landeswahlleiterin, etwa zwei Wochen vor der Wahl erreichen die Vorbereitungen ihren vorläufigen Höhepunkt. In umfangreichen Tests werden in den verbleibenden Tagen bis zum Wahlabend die Erfassungs- und Präsentationsprogramme, die Module zur Sitzberechnung, die Prozesssteuerung sowie die Druckausgaben auf ihre fachliche und formale Korrektheit geprüft. Durch die räumliche Nähe zum Auftraggeber, der Landeswahlleiterin, ergeben sich mitunter noch Ad-hoc-Anforderungen, die in dieser Zeit ebenfalls erstellt und in das Wahlverfahren eingebaut werden müssen.

Mit dem Umzug können auch erstmals die Schaltungen aller Leitungen erprobt werden. Aus dem lokalen Wahlnetz des LDS NRW werden die Daten über das Landesverwaltungsnetz (LVN) zu den Servern für die Präsentation im Web oder für den Ergebnisabruf in den Landtag übertragen. Bei Wahlen auf Bundesebene müssen die Wahlergebnisse außerdem an den Bundeswahlleiter übermittelt werden. Um eine größtmögliche Ausfallsicherheit in der Netztechnik sicherzustellen, werden gesonderte und unabhängige Leitungen angemietet, auf die im Störfall umgeschaltet werden kann. Die Leistungsbreite und -stabilität des LVN hat sich

in den vergangenen Jahren aufs Beste bewährt, so dass nie auf die Backup-Lösungen zurückgegriffen werden musste.

Zur Landtagswahl bauen die Medien im Landtagsgebäude ihre Wahlstudios auf, aus denen live am Wahlabend gesendet wird. In der Bürgerhalle findet eine große „Wahlparty“ mit rund 2 000 Besuchern statt. Auch die Landeswahlleiterin zeigt hier Präsenz: Im Sitzungssaal der „Landespressekonferenz“ steht ein Team der Landeswahlleiterin und des LDS NRW bereit, um die Besucher umfangreich und aktuell mit Wahldaten zu versorgen. In der Bürgerhalle präsentiert ein Großbildschirm aktuelle Informationen. Individuelle Datenabrufe können an dezentral aufgestellten Kiosksystemen von den Besuchern selbst vorgenommen werden. Diese Systeme arbeiten ähnlich wie ein Fahrkartenautomat: Mit dem Finger auf der berührungssensitiven Scheibe steuert der Anwender die Handhabung, selektiert einen Wahlkreis, ruft den Druck auf. Natürlich stammt auch diese Anwendung vom Wahlteam des LDS NRW.



Abb. 1: Das „Kiosksystem“ mit dem berührungssensitiven Touchscreen

1) Siehe Beitrag in dieser Ausgabe der LDVZ-Nachrichten „Landtagswahl 2005 – Interaktive Kartendarstellung der Wahlergebnisse“ von Anja Neumann, Georg Stahl und Christoph Rath.

Abb. 2: Die Startseite des „Kiosksystems“

Der Wahlabend

Rund eine Stunde vor Schließung der Wahllokale wird das Wahlnetz hochgefahren und einem letzten Test unterzogen. Wie alle Wahlbürgerinnen und -bürger schart sich auch das Wahlteam um 18 Uhr um einen Fernseher, um die ersten Trendmeldungen der Fernsehanstalten zu verfolgen. Diese Vier-



Abb. 3: Eine Mitarbeiterin des LDS NRW bei der telefonischen Erfassung eines Wahlkreises

telstunde ist nicht nur dem politischen Interesse geschuldet. Die Trends haben in der letzten Zeit eine so gute Qualität erreicht, dass sie zur Skalierung der Plausibilitätsprüfungen verwendet werden können: Hat z. B. die Partei A in der Vergangenheit i. d. R.

18 % der Stimmen erhalten, werden die Erfassungsprogramme mit Unter- und Obergrenzen (z. B. 12 % – 24 %) hinterlegt, deren Unter- oder Überschreitung eine Warnmeldung erzeugt. Damit können u. U. Fehleingaben oder Übermittlungsfehler frühzeitig aufgedeckt werden. Vermeldet nun der Trend einen Einbruch für die Partei A auf z. B. 7 %, so können die Grenzen umgehend auf die neue Situation angepasst werden, um die Warnmeldungen zu reduzieren.

Um die eingehenden Ergebnisse aus den Wahlkreisen in die Wahldatenbank aufnehmen zu können – gemeldet wird von den Kreiswahlleiterinnen und Kreiswahlleitern per Telefon und Fax, werden acht Erfassungsarbeitsplätze eingerichtet, an denen entweder die auf Papier vorliegenden Daten erfasst oder im fernmündlichen Dialog direkt in die An-

wendung aufgenommen werden. Die inzwischen überwiegend eingesetzte Dialogkomponente hat den großen Vorteil, dass aufgrund der implementierten Plausibilitätsprüfungen Unstimmigkeiten in den Ergebnissen sofort bemerkt und korrigiert sowie Rückfragen reduziert werden können.

Es mutet vielleicht anachronistisch an, dass herkömmliche Medien wie Fax und Telefon und nicht Web-basierte Übertragungswege für die Übermittlung der aktuellen Ergebnisse Verwendung finden. Die Vorteile liegen indessen auf der Hand: Bei einer Landtagswahl sind insgesamt lediglich rund 5 000 Datenwerte zu übermitteln, bei Kommunalwahlen und der Bundestagswahl sind es nur unwesentlich mehr. Eine elektronische Datenübermittlung müsste aus Sicherheitsgründen aufwändig über diverse Schutzmechanismen (Firewalls, Zuteilung von Übertragungsrechten usw.) angebunden werden, was angesichts des Datenvolumens kaum wirtschaftlich erscheint. Außerdem funktioniert das Telefon im Prinzip immer und ist für moderne Hacker-Angriffe uninteressant.

Bei der diesjährigen Landtagswahl am 22. 5. 2005 ging das erste Wahlkreisergebnis um 19:04 Uhr ein. Der Wahlkreis 76-Bottrop konnte damit die

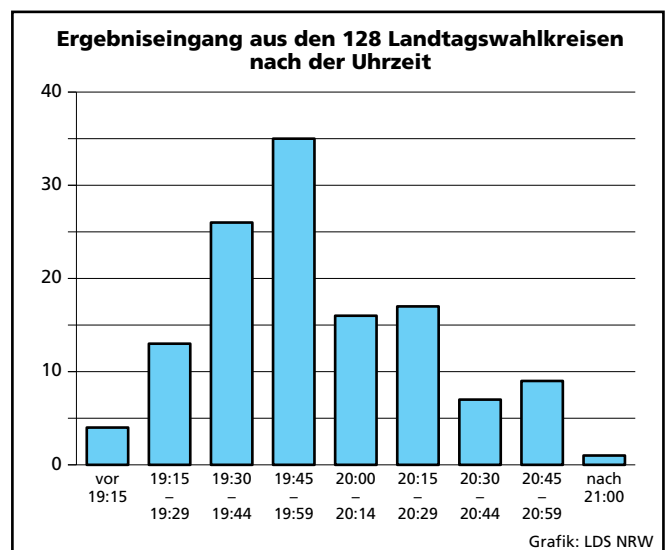


Abb. 4: Der zeitliche Eingang der Wahlkreisergebnisse bei der Landtagswahl am 22. Mai 2005

„Sieger“ vieler zurückliegender Wahlen, Düsseldorf I-IV, um zwei Minuten schlagen. Obwohl die Stadt Köln flächendeckend Wahlmaschinen einsetzt, kommt die letzte Meldung um 21:01 Uhr aus dem Wahlkreis 17-Köln V.

Abbildung 4 verdeutlicht den zeitlichen Verlauf der Eingänge. Bis 20 Uhr sind bereits mehr als die Hälfte der Ergebnisse aller Wahlkreise eingegangen. Das Maximum ist um 19:54 zu verzeichnen, als fast zeitgleich fünf Erfasserinnen die Daten entgegennehmen müssen. Dank der hochwertigen Arbeit in den Wahllokalen und Wahlämtern der Kommunen und Kreise und der stetigen Rückfragen bei echten oder scheinbaren Unstimmigkeiten ist das Material praktisch fehlerfrei, lediglich zwei Korrekturen werden am Wahlabend nachgetragen. Exakt 41 Minuten nach dem Eingang des letzten Wahlkreisergebnisses sind das Landesergebnis und die Sitzverteilung um 21:42 Uhr errechnet, geprüft und werden durch die Landeswahlleiterin im Landtag veröffentlicht. Das Wahlteam erstellt sogleich die Druckvorlagen, damit die Hausdruckerei des LDS NRW in einer kompakten Nachtschicht die gewünschten 1 300 Exemplare des 200 Seiten starken Heftes drucken, binden und verteilen kann, so dass pünktlich zu Dienstbeginn im Landtag, bei den politischen Parteien, in den Ressorts und weiteren Stellen der Landesverwaltung die druckfrischen Exemplare im Posteingang liegen.

Neben dieser konventionellen Ergebnisverbreitung stellt das Internet heute sicher das bedeutsamere Kommunikationsmedium dar. Mit der Einspeicherung des jeweils eingegangenen Wahlkreises in die Wahl-Datenbank wird automatisch eine HTML-Seite generiert und auf den Webserver übertragen. Dort führt eine geordnete Liste alle Wahlkreise auf, die von einem Dämon überwacht und aktualisiert wird, so dass jederzeit an einem grünen Haken erkenntlich ist, welche Wahlkreise inzwischen vorliegen. Nach Abschluss der

Wahl werden Dateien in verschiedenen Formaten (PDF, CSV) mit allen Wahlergebnissen im Web zum Download bereitgestellt.

Mit einem Mausklick kann der Anwender im Web auch auf gesonderte Seiten navigieren, auf denen die Wahlergebnisse kartografisch aufbereitet werden. Im Zentrum findet sich eine Karte mit den Wahlkreisen des Landes, die nach verschiedenen Merkmalen eingefärbt

werden kann. Es wird angezeigt, welcher Bewerber das Direktmandat gewonnen hat, wie die Wahlbeteiligung in den Wahlkreisen war, welchen Stimmenanteil die SPD, die CDU, die GRÜNEN oder die FDP errungen haben. Um sich einen Detailblick zu verschaffen, kann der Anwender in die Karte hineinzoomen oder über Auswahlboxen direkt Kommunen ansteuern. Die Karte wird nach der Auswahl automatisch auf den Wahlkreis ausge-



Abb. 5: Landeswahlleiterin Helga Block (vorne links) im Informationsraum im Landtag



Abb. 6: Ein Besucher fragt nach aktuellen Ergebnissen.

richtet, zu dem die gesuchte Kommune gehört. Selbstverständlich können die Karten gedruckt werden, ebenso ist ein Umschalten auf die tabellarische Darstellung der Wahlergebnisse jederzeit möglich. Dieses GIS-basierte Web-Modul ist selbstverständlich intuitiv zu handhaben.

Wie überaus intensiv das Wahlangebot www.wahlen.nrw.de im Web inzwischen genutzt wird, verdeutlichen die Zugriffszahlen: Am Wahlabend wurden mehr als 1,9 Millionen Mal die Seiten aufgerufen, auch am Folgetag war der Zugriff mit 800 000 Treffern noch beachtlich. Zwischen 19 und 22 Uhr lag die Nutzung bei mehr als 300 000 Zugriffen in der Stunde, zwischen 20 und 21 Uhr wurde eine Spitze mit 466 588 Anfragen ausgelöst. Das gesamte übertragene Datenvolumen beträgt über 50 Gigabytes. Die Zugriffe auf die GIS-Komponente lassen sich nicht so exakt ermitteln, da die Seitenanfragen z. T. über Frames laufen und dadurch die Zähler überhöhte Werte anzeigen, sie dürfte aber auch knapp die halbe Million mit einem Übertragungsvolumen vom 26 Gigabytes erreicht haben. Damit wurde eine Last auf den Servern ausgelöst, die mehr als doppelt so hoch ist als die Maxima an durchschnittlichen Arbeitstagen. Diese Spitzenbelastung wurde dank der Zugangsbandbreite in das Web von 155 MBit und der Ausstattung der GIS-Server mit einem Loadbalancer spielend gemeistert; die Wahlen stellen somit für das LDS NRW nebenbei eine verlässliche Möglichkeit zur Performanzmessung der Infrastruktur unter realen Bedingungen dar.

Spannend ist ein Blick auf die Domänen, aus denen das Web-Angebot aufgerufen worden ist, lässt sich daraus doch eine Aussage über das überregionale Interesse ableiten. Natürlich kommt der überwiegende Teil (80 %) der Anfragen aus Deutschland (.de und .net) und dem benachbarten Ausland (12 %). Der verbleibende Rest verteilt sich offenkundig auf die ganze Welt:

Zum Surfen in den Ergebnissen der Landtagswahl 2005 haben sich Anwender z. B. in Tansania, Tuvalu, Peru, Kuba, Brunei, Laos, Mauritius, Mozambique und auf der Weihnachtsinsel vor ihren PC gesetzt!

Nach der Wahl

Zunächst einmal gilt es, möglichst schnell die für die Wahl okkupierten Räume freizugeben. So muss noch in der Wahlnacht der Raum der Landespresskonferenz geräumt werden; der Raum im Innenministerium steht schon am Nachmittag nach der Wahl wieder als Besprechungsraum zur Verfügung. Die Wahlarbeiten selbst sind allerdings keinesfalls beendet. Am folgenden Wochenende werden die endgültigen Ergebnisse, die im Laufe der Woche nach den Detailprüfungen der Kreiswahlleiter ermittelt worden sind, sowie die Ergebnisse auf Gemeindeebene erfasst, geprüft, miteinander verglichen, ins Web gestellt und gedruckt. Es wird praktisch der Wahlabend mit einigen Erweiterungen wiederholt, allerdings ohne dessen zeitlichen Druck. Einige Wochen später werden die Ergebnisse der repräsentativen Wahlstatistik aufgenommen und ebenfalls veröffentlicht, erst dann kehrt im Wahlbereich wieder Ruhe ein.

Sicher gibt es ungleich umfangreichere DV-Verfahren im LDS NRW. Gleichwohl stellt die IT-Unterstützung der Landeswahlleiterin bei allen Wahlen, insbesondere aber bei Landtagswahlen, eine besondere Herausforderung an die Mitarbeiterinnen und Mitarbeiter des LDS NRW dar, gilt es doch sowohl eine komplexe IT-Dienstleistung punktgenau zu erbringen als auch absolut verlässliche Ergebnisse in der Wahlnacht zu ermitteln. Ermöglicht wird dies seit über zwei Jahrzehnten durch eine ausgezeichnete und engagierte, referatsübergreifende Zusammenarbeit der Hauslogistik und Druckerei, des Referates für Öffentlichkeitsarbeit, der LAN-, WAN- und

Web-Administratoren sowie des IT-Servicecenters, vor allem aber durch die Kooperation zwischen dem für die Verfahrensentwicklung verantwortlichen IT-Referat und dem für die Wahlen zuständigen Fachreferat.

Ein Exkurs der Technik

Die gesamte Wahanwendung beruht auf PC-Technik. Zum Einsatz kommen sowohl herkömmliche PC und Laptops, um die Transportaufwände zu minimieren. Als Datenbanksystem wird MS-SQL in der Version 2000 verwendet, als Betriebssystem kommt bei den Servern Windows 2003, bei den Clients Windows XP zum Einsatz, die Programmierung erfolgt ausschließlich mit Delphi 7. Die GIS-Komponente wird auf der Basis von ARCIMS mit Java realisiert. Zwischenzeitliche Überlegungen, das Verfahren auf eine andere Plattform zu migrieren und auf Open-Source-Software umzustellen, wurden verworfen, da weder Geschwindigkeitsvorteile noch finanzielle Einsparungen zu erwarten sind.



*Dr. Thomas Pricking
Tel.: 0211 9449-5058
E-Mail: thomas.pricking@lds.nrw.de*

Landtagswahl 2005

Interaktive Kartendarstellung der Wahlergebnisse

Bei der Landtagswahl wurde die bewährte Präsentation der Ergebnisse im Internet erstmals um eine interaktive Kartendarstellung ergänzt. In sechs inhaltlichen Ebenen werden die Gewinner/-innen der Direktmandate, die Wahlbeteiligung und die prozentualen Stimmenanteile der vier großen Parteien dargestellt. Mittels eines Klicks in die Karte kann das detaillierte Ergebnis eines Wahlkreises aufgerufen werden. Über verschiedene Funktionen, z. B. die Auswahl von Regierungsbezirk, Kreis oder Gemeinde, können einzelne Landesteile genauer betrachtet werden. Eine Funktion zur Druckausgabe rundet die Anwendung ab. Das Layout ist auf den Internetauftritt der Landeswahlleiterin, bzw. des Innenministeriums abgestimmt und passt sich somit nahtlos in die klassische Präsentation der Wahlergebnisse ein.

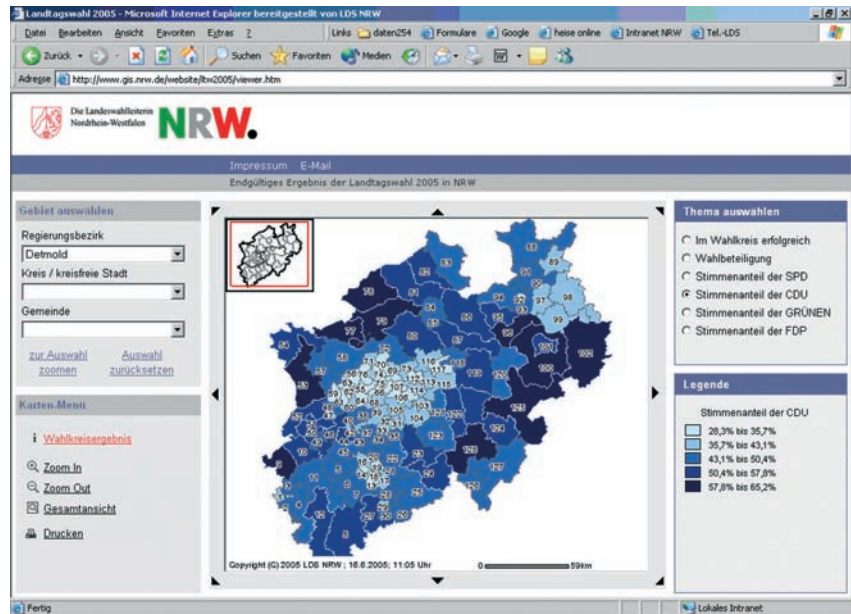


Abb. 1: Applikation hier mit Wahlergebnis der CDU

Anwendung

Die Anwendung strukturiert sich klar in fünf Bereiche. Neben dem eigentlichen Kartenbereich sind dies:

Gebiet auswählen

Regierungsbezirk
Detmold

Kreis / kreisfreie Stadt
Gütersloh

Gemeinde

[zur Auswahl zoomen](#) [Auswahl zurücksetzen](#)

Gebiet auswählen

Über Klapplisten können Regierungsbezirk, Kreis oder Gemeinde ausgewählt werden. Mit dieser Funktion kann der Kartenausschnitt einfach auf eine bestimmte Region fokussiert werden.

Karten-Menü

[Wahlkreisergebnis](#)

[Zoom In](#)

[Zoom Out](#)

[Gesamtansicht](#)

[Drucken](#)

Karten-Menü

Neben weiteren Funktionen zur Auswahl des Kartenausschnittes – Zoom In, Zoom Out, Gesamtansicht – kann hier die Kartenausgabe im druckerfreundlichen Format aufgerufen werden. Durch Aktivierung des Werkzeugs

Wahlkreisergebnis wird durch einen Klick in die Karte das detaillierte Wahlkreisergebnis in einem Pop-up-Fenster angezeigt. Dies erfolgt durch Verlinkung auf das klassische Angebot.

Thema auswählen

☒ Im Wahlkreis erfolgreich

☐ Wahlbeteiligung

☐ Stimmenanteil der SPD

☐ Stimmenanteil der CDU

☐ Stimmenanteil der GRÜNEN

☐ Stimmenanteil der FDP

Thema auswählen

In diesem Frame werden die unterschiedlichen inhaltlichen Kartenebenen angeboten. Da immer nur eine Ebene sinnvoll dargestellt werden kann, wird die Auswahl in Form von RadioButtons angeboten.

Legende

Stimmenanteil der GRÜNEN

3,0% bis 6,1%

6,1% bis 9,2%

9,2% bis 12,4%

12,4% bis 15,5%

15,5% bis 18,6%

Legende

In der Legende wird die jeweils gewählte Kartenebene (Thema) erläutert. Bis auf die Ebene für den Gewinn des Direktmandats stellen alle anderen Ebenen prozentuale Anteile

dar. Die Klassierung erfolgt jeweils in fünf Klassen. Ausgehend von eng gesetzten Startwerten, die sich an den letzten Prognosen orientierten, wurden die Werte durch die einlaufenden Wahlergebnisse bestimmt.

Technik

Die technischen Anforderungen wurden vor allem durch eine kaum kalkulierbare hohe Nachfrage in der Wahlnacht bestimmt. Für besondere Spannung sorgte zudem die Tatsache, dass es keine Möglichkeit zur Nachbesserung gab. Die Anwendung musste in der Wahlnacht stabil und performant laufen. Aus diesen Randbedingungen wurde folgende Vorgehensweise entwickelt:

- Kalkulation der zu erwartenden Last, primär unter Beachtung der Zugriffszahlen bei den letzten Kommunalwahlen
- Konzeption der Anwendungsarchitektur mit dem Fokus auf minimale Server- und Netzlast
- Entwicklung eines Prototypen
- Durchführung eines Lasttest zur Ermittlung der benötigten Ressourcen
- Konzeption der Infrastruktur mit redundanter Auslegung von Servern und Netzkomponenten
- Anwendungsentwicklung und Dokumentation
- Aufbau der Infrastruktur
- Lasttest in Produktionsumgebung

Infrastruktur

Als Ergebnis dieser Überlegungen wurde eine Konfiguration mit 10 GIS-Servern gewählt, um in der Wahlnacht auch bei sehr hoher Last die Nachfrage bedienen zu können. Da die Server nur eine kurze Zeit für die Landtagswahl gebunden waren und inzwischen anderen Projekten zugeführt wurden, blieben die Aufwände überschaubar. Ein Loadbalancer verteilte die Sessions nahezu gleichmäßig auf die GIS-Server.

Anwendungsarchitektur Server

Der GIS-Server gliedert sich in mehrere Komponenten auf, die jeweils ein eigenes Software-Produkt darstellen.

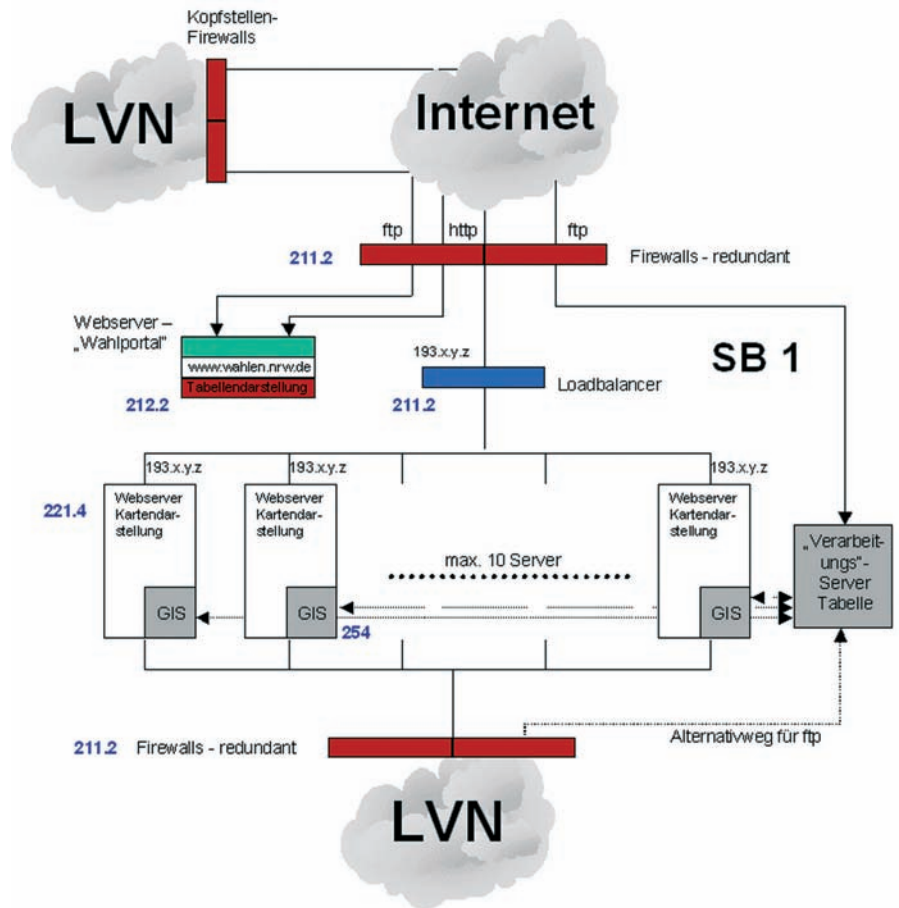


Abb. 6: Infrastruktur

Die Module der Komponenten ermöglichen deren Kommunikation untereinander. Die folgende Abbildung stellt die Architektur mit den Modulen dar, die in der Web-Anwendung zur Anwendung kommen. Sie beschränkt sich auf die Darstellung eines einzelnen GIS-Servers.

Vom Client werden die Karten über so genannte AXL-Requests (AXL steht für ArcXML, eine auf XML-basierende von der Firma ESRI entwickelte Beschreibungssprache) vom Server angefordert. Der GIS-Server erzeugt die Karte und sendet die URL unter der die Karte angefragt werden kann, an den Client. Der Client kann die Karte dann abholen und darstellen.

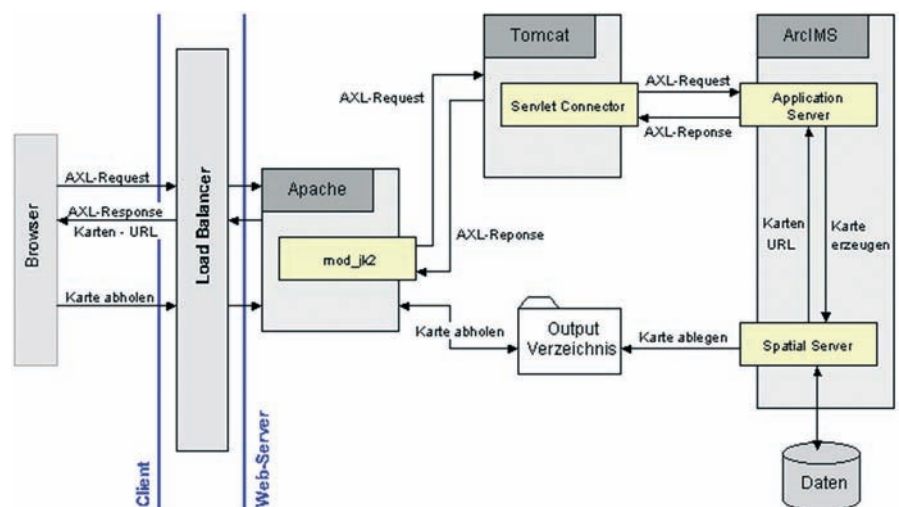


Abb. 7: Anwendungsarchitektur Server

Anwendungsarchitektur Client

Der Client gliedert sich in mehrere Frame-Komponenten auf. Jede Komponente greift auf Module in Form von HTML-Seiten und/oder JavaScript-Dateien zu. Die folgende Abbildung stellt die Architektur mit den wesentlichen Komponenten und den dazugehörigen wichtigsten Modulen dar.

Die funktionalen Elemente wurden bereits oben beschrieben. Besonderes Augenmerk wurde noch auf die Ablauffähigkeit in den verschiedensten Browser-Umgebungen gelegt. Dazu werden verschiedene nutzerspezifische Browser-Einstellungen, wie JavaScript und Pop-up-Unterstützung beim Start der Anwendung getestet. Sofern erforderlich, wird dem Nutzer ein Hinweis zur Änderung dieser Einstellungen und eine Beschreibung zur Durchführung dieser Änderungen gegeben.

Anwendungsarchitektur Verarbeitungsserver

Der Workflow zur schnellen Aktualisierung der Ergebnisdarstellung stellt sicher eine besondere Herausforderung dar. Technisch orientiert sich die Lösung hierbei am bisherigen Angebot. Die Ergebnisse der Wahlkreise werden mittels ftp auf zwei Verarbeitungsservern (redundante Auslegung) abgelegt. Minütlich wurde auf diesen Servern in der Wahlnacht abgefragt, ob neue Wahlergebnisse eingegangen sind. Sofern dies der Fall war, wurden die Daten für die Karten und die Legendengrafiken neu berechnet und auf die Kartenserver transferiert. Klassisches Angebot und Kartendarstellung konnten so praktisch synchron gehalten werden.

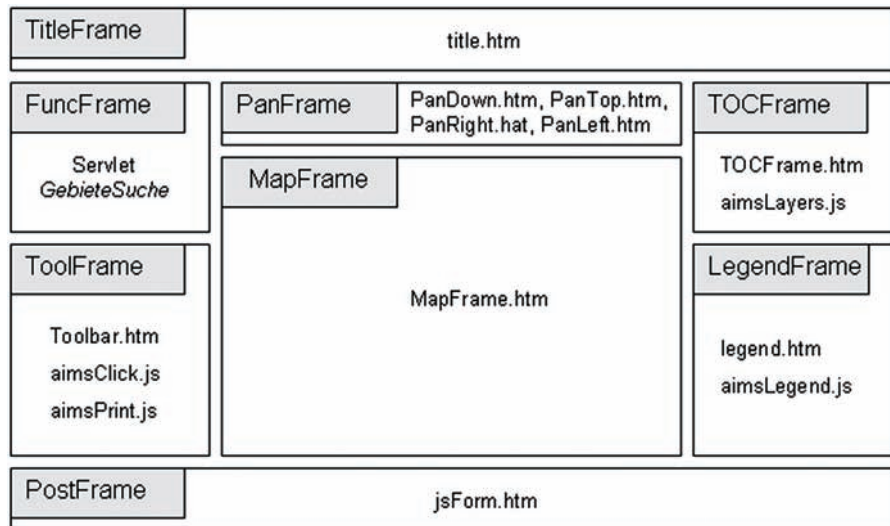


Abb. 8: Anwendungsarchitektur Client

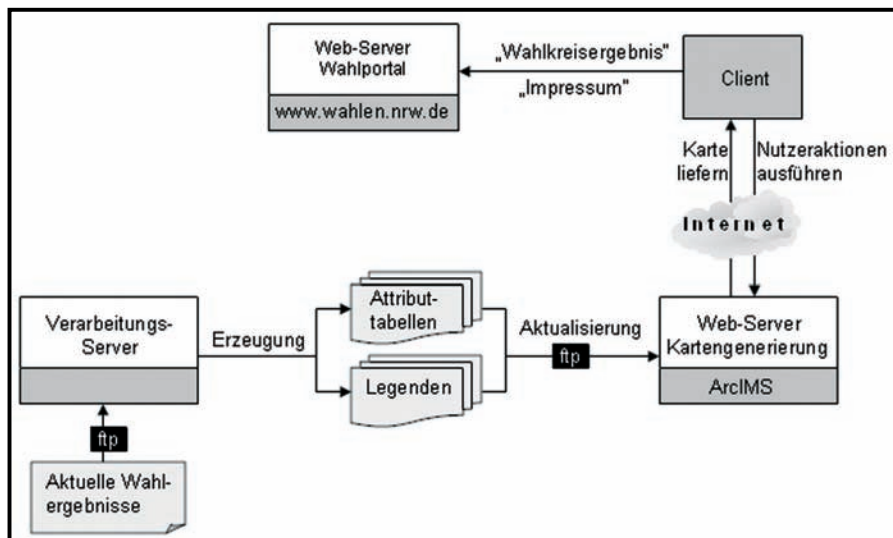


Abb. 9: Anwendungsarchitektur Verarbeitungsserver

Fazit

Die Ergänzung der Internetpräsenz der Landeswahlleiterin NRW um eine interaktive Kartendarstellung kann als rundum positiv bewertet werden. Besonders erfreulich war aus Sicht des Geoinformationszentrums die sehr effektive und konstruktive Zusammenarbeit beteiligter Referate im LDS NRW.

Gesamtangebot Landtagswahl 2005:
<http://www.wahlen.lds.nrw.de/landtagswahlen/2005/index.html>

Interaktive Kartendarstellung der endgültigen Wahlergebnisse:
<http://www.gis.nrw.de/website/ltw2005>



Anja Neumann
 Tel.: 0211 9449-6325
 E-Mail: anja.neumann@lds.nrw.de



Georg Stahl
 Tel.: 0211 9449-6315
 E-Mail: georg.stahl@lds.nrw.de



Christoph Rath
 Tel.: 0211 9449-6318
 E-Mail: christoph.rath@lds.nrw.de

Aktuelle Trends in der Spracherkennung

Der Einsatz von Spracherkennungssoftware hat sich im Alltag in vielfältigen Einsatzgebieten bereits fest etabliert. Ein Anruf bei einer Bank oder einer Telefonauskunft landet mit hoher Wahrscheinlichkeit zunächst in einem Sprachdialogsystem (Voice Portal). Ein solches sprachgesteuertes System führt den Anrufer oftmals bis zur abschließenden Bearbeitung seines Anliegens, zumindest jedoch kann es ihn mit der zuständigen Stelle verbinden. Die Sprachwahlmöglichkeit bei modernen Mobiltelefonen ist heute quasi Standard: Nach entsprechender Konfiguration nennt der Anrufer seinen gewünschten Gesprächspartner, das Handy wählt die zugeordnete Nummer. Auch in vielen weiteren Bereichen wird Spracherkennung erfolgreich eingesetzt.

Aktuelle Einsatzgebiete

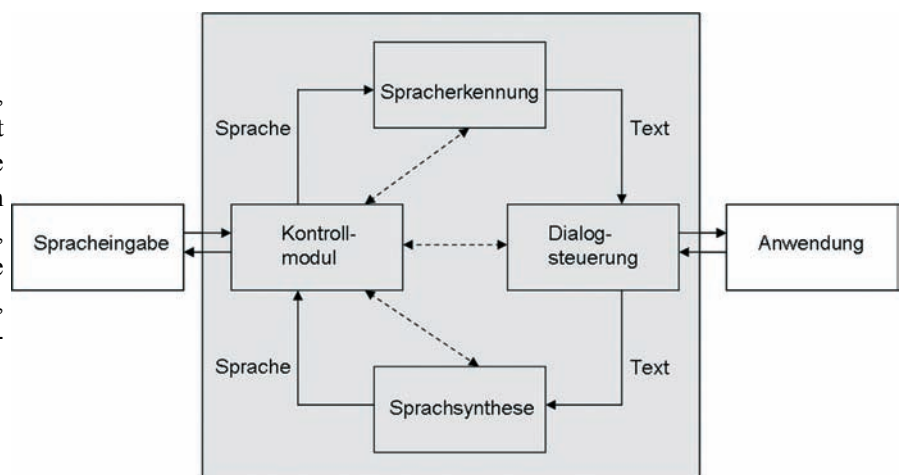
Spracherkennung gliedert sich in verschiedene Teilbereiche mit unterschiedlichen Einsatzmöglichkeiten:

- Steuerung von Maschinen
Die Befehle sind sprecherunabhängig, eine geringe Menge an Befehlen ist hierbei ausreichend. Beispiele sind die Steuerung von Haushaltsgeräten (auch per Telefon) oder Werkzeugmaschinen, mobile Kommunikationsgeräte wie Handy, Smartphone oder auch PDA, Systeme im KFZ wie Radio, Navigationssystem etc.
- Befehlseingaben an Computersysteme
Diese erfolgen primär über Telefon. Auch diese Systeme arbeiten sprecherunabhängig und verstehen nur einen relativ geringen Wortschatz (Schlüsselworte). Beispiele sind Datenbankrecherchen mittels der schon genannten Sprachdialogsysteme, insbesondere verschiedene Arten von Buchungssystemen, Auskunftssystemen, Bankingsystemen etc.
- Einsatz von Spracherkennungssoftware
Auf PCs werden diktierete Texte erfasst, Programme werden per Sprache gesteuert. Bei dieser Nutzungsvariante werden generell andere Schwerpunkte gesetzt: Ein System zur Spracherkennung, welches komplette gesprochene

ne Sätze erfassen soll, ist (noch) sprecherabhängig und setzt somit einen gewissen Trainingsaufwand voraus. Hierbei „lernt“ das Programm die Besonderheiten der Stimme des Diktierenden. Da die Diktatinhalte keinen speziellen Einschränkungen unterliegen sollen wird mit sehr großen Wortschätzen (Vokabularen) gearbeitet. Einen Sonderfall stellt noch die Nutzung spezieller Fachwortschätze dar, z. B. für die Wirtschaft, für Juristen oder Fachärzte. Hierbei handelt es sich allgemein um Erweiterungen der Standardwortschätze für bestimmte Fachbereiche.

Spracherkennungssoftware wird nicht nur für allgemeine Diktate und zur Bedienung von PCs per Sprachbefehlen verwendet, sondern oftmals auch im Rahmen von Sprachlern- und Übersetzungsprogrammen.

Die Eingabemöglichkeit für Sprachbefehle ist heute z. T. bereits in Betriebssystemoberflächen oder Anwendungen integriert (Windows XP, neuester Opera-Browser etc.).



Funktionsweise eines Sprachdialogsystems

Ein kurzer Überblick

Bereits seit den frühen 1960er-Jahren wird an Systemen zur Erkennung der menschlichen Sprache geforscht und entwickelt. Beteiligt waren dabei verschiedene Labore, größtenteils in den USA. Bis in die 1980er-Jahre hinein ging es dabei nur um das Erkennen von maximal einigen 100 Einzelworten. Bedeutende Technologiemeilensteine setzte dabei oftmals die Fa. IBM. Diese entwickelte beispielsweise be-

reits 1962 ein Gerät zur Sprachausgabe. Im Jahr 1986 wurde mit dem „Tangora 4“ der Prototyp eines echten Spracherkennungssystems (Soft- und Hardware) vorgestellt, der bereits viele Problembereiche der Spracherkennung zu lösen versuchte. So verfügte er über Möglichkeiten zur Kontextanalyse, Unterscheidung von Homofonen und Trigrammstatistiken.

Im Jahr 1996 stellte IBM eine erste hardwareunabhängige Spracherkennungslösung für Windows 95® vor. Erstmals konnte nun mit Standard-soundkarten unter einer weltweit verbreiteten Betriebssystemumgebung gearbeitet werden.

Der für die Spracherkennung entscheidende Durchbruch erfolgte dann im Jahr 2001: Die Vorstellung der kontinuierlichen Spracherkennung.

Unterschiedliche Diktierarten

Die bis zu diesem Zeitpunkt übliche „diskrete“ Diktierweise machte es erforderlich, jedes Wort für sich, gefolgt von einer deutlichen Sprechpause, auszusprechen. Dies erforderte eine hohe Konzentration und Diktierdisziplin und hörte sich ziemlich unnatürlich an. Jedoch erleichterte sie dem Spracherkennungssystem die Arbeit: Da die Wörter klar abgegrenzt diktiert wurden, waren ausgefeilte Routinen zur Zerlegung eines Satzes in Einzelworte nicht erforderlich.

Bei der ab dem Jahr 2001 verfügbaren „kontinuierlichen“ Diktierweise wird wie bei einem „normalen“ Gespräch diktiert. Die Worte eines Satzes sind nahezu ohne hörbare Laut- und Wortgrenzen aneinander gereiht. Das Spracherkennungssystem muss hierbei

den kontinuierlichen Redefluss zunächst strukturieren und in Einzelworte zerlegen, um ihn anschließend analysieren zu können. Hierbei kommen mehrere Techniken zum Einsatz.

Wie wird Sprache erkannt?

Die Berechnungsprozesse, die schließlich zum Erkennen von Sprache führen, sind überaus komplex. Die wichtigsten Schritte hierbei sind:

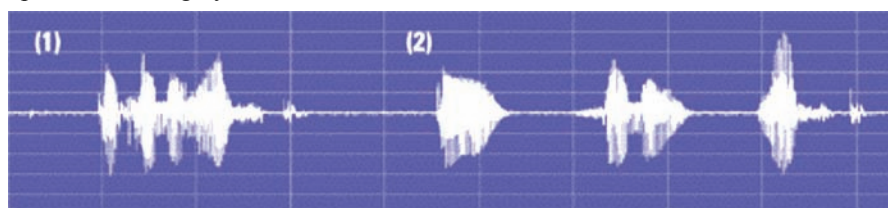
- 1) Der Sprecher diktiert in ein (analoges) Mikrofon (Headset oder Standmikrofon).
- 2) Das analoge Sprachsignal wird digitalisiert. Dieser Vorgang erfolgt normalerweise in der Soundkarte des eingesetzten Rechners. Bei Verwendung eines speziellen Mikrofons (z. B. USB-Headset mit eingebautem Rauschfiltern zur Klangverbesserung) findet die Digitalisierung direkt in diesem statt. Diese Lösung liefert erfahrungsgemäß meist das beste Sprachsignal.
- 3) Das digitalisierte Sprachsignal wird in sog. „Foneme“ zerlegt. Ein Fonem ist im Sinne der Spracherkennung die kleinste bedeutungsunterscheidende lautsprachliche Einheit (also quasi das „Bit“ des Diktats). Die deutsche Sprache z. B. nutzt etwa 40 derartige Foneme.
- 4) Die Foneme werden mit im System hinterlegten Referenzmustern verglichen, um sie eindeutig identifizieren und zuordnen zu können.
- 5) Um aus der Musterfolge nun einzelne Worte herauslesen zu können gibt es verschiedene Berechnungsmodelle für die akustische Modellierung, die nun zur Anwendung kommen können. Das bekannteste Verfahren ist das sog. „Hidden-Markov-Modell“, es beschreibt die Wahrchein-

lichkeiten der Abfolge bestimmter Muster. Als Abschluss der Berechnungen hat das System die einzelnen Worte – hoffentlich richtig – erkannt.

- 6) Es folgt die Anwendung des sog. Sprachmodells. Es ermöglicht einen Wortvergleich anhand einer Trigramm- bzw. Bigrammstatistik (3 bzw. 2 aufeinander folgende Worte). Ziel ist es, die Wahrscheinlichkeit zu ermitteln, dass auf ein erkanntes Wort ein anderes sinnvoll passendes Wort folgt.
- 7) In der anschließenden Suche wird nun in einem rechenintensiven Prozess aus allen vorliegenden Informationen die gesprochene Wortfolge ermittelt.
- 8) Der erkannte Text wird abschließend auf dem Bildschirm ausgegeben.

Da die genannten Prozesse sämtlich in Echtzeit ablaufen, sind Programme zur Spracherkennung immer sehr rechenintensiv und stellen bei anzustrebender hoher Genauigkeit sehr hohe Anforderungen an die eingesetzte Hardware (CPU-Leistung, Hauptspeicher). Ein ausreichender Hauptspeicher ist doppelt wichtig, da die verwendeten Sprachmodellldaten dort permanent vorgehalten werden müssen, um den sichtbaren Nachlauf des Systems (die Zeit vom Diktieren eines Wortes bis zur Anzeige auf dem Bildschirm) so gering wie möglich zu halten. Ein deutlicher Nachlauf erschwert das Diktieren erheblich, insbesondere wenn konzeptionell an vorhandenem Text gearbeitet werden soll.

Der erstellte Text bedarf nun in jedem Fall noch einer Korrektur bzw. Überarbeitung durch den Diktierenden. Die einzelnen Worte sind zwar, da sie aus einem Vokabular entnommen wurden, fast immer richtig geschrieben, können jedoch vom Sinn her völlig falsch sein. Die Erkennungsleistung eines Systems liegt heute zwar – nach einer gewissen Trainingsphase – meist bei 95 – 98 %, jedoch gibt es beim Diktieren eine Reihe von Problemen, die der Software die Arbeit teils erheblich erschweren.



(1) kontinuierliche Sprechweise

(2) diskrete Sprechweise

Sprecherunabhängige Systeme

Ein sprecherunabhängiges System kann von jeder Person unmittelbar genutzt werden, ohne es vorher trainieren zu müssen. Hierbei ist zu differenzieren, ob sich die Sprecherunabhängigkeit nur auf den Wortschatz oder auf das komplette System (wie bei Voice Portalen) bezieht. Obwohl ein Anpassungstraining in diesen Fällen nicht mehr notwendig (bzw. möglich) ist, ist es aufgrund der individuellen Sprachgewohnheiten bei einem trainierbaren System nach wie vor empfehlenswert, da es die Erkennungsgenauigkeit in jedem Fall noch erhöht.

Probleme bei der Erkennung von Sprache

Die bei der Nutzung von Spracherkennungslösungen auftretenden Problemfälle können vielfältiger Natur sein. Man unterscheidet hier Ursachen technischer, fonetischer oder linguistischer Art.

Zu den technischen Problemen zählen vorrangig die Eigenschaften von Mikrofon und Soundkarte (Qualität der Hardware), Umgebungsgeräusche (Großraumbüro, dauerndes Telefonklingeln etc.), ungenügende Hardwareleistung (CPU, Hauptspeicher, Festplatte). Bei Nutzung eines Telefons wirkt sich natürlich auch die Qualität der Verbindung aus.

Zu den Problemen fonetischer Art zählen insbesondere Variationen in der Aussprache der/des Diktierenden (Konzentration, Geschwindigkeit, Betonung, Dialekt, Gesundheitszustand).

Die Reihe möglicher linguistischer Probleme ist ebenfalls vielfältig. So können beim Diktieren Schwierigkeiten bei der korrekten Erkennung von Homophonen (viel – fiel), Komposita, Groß-/Kleinschreibung, Derivationen, Flexionen und der Satzstellung auftreten. Außerdem erkennt kein PC die

Gestik und Mimik eines Sprechers. Diese zusätzlichen Informationen ermöglichen menschlichen Gesprächspartnern oftmals erst die Erkennung des Sinns einer Aussage (und damit der verwendeten Worte).

Heutiger Stand

Auf dem Markt wird heute eine ganze Reihe von Produkten zur Spracherkennung angeboten. Die Produkte sind jedoch zumeist keine Eigenentwicklungen, sondern oftmals (teils auch ältere) OEM-Versionen der Produkte der Firmen IBM und ScanSoft (Dragon). Beide Hersteller entwickeln ihre heute nahezu gleichwertigen Produkte ständig weiter.

Die aktuell marktverfügbaren Versionen sind IBM ViaVoice 10 und Dragon NaturallySpeaking 8. Beide Produkte werden in verschiedenen Ausführungen (und somit Preiskategorien) für Einsteiger, fortgeschrittene, professionelle sowie mobile Nutzer/-innen (im Bundle mit einem digitalen Diktiergerät) vermarktet. Darüber hinaus werden spezielle Fachvokabulare für die Bereiche Wirtschaft, Recht und Medizin angeboten.

Ein Großteil der marktgängigen PC-gestützten Spracherkennungslösungen wird in den letztgenannten Bereichen (Rechtsanwaltskanzleien, Arztpraxen, Krankenhäusern etc.) eingesetzt.

Weitere ehemalige Anbieter eigener Spracherkennungslösungen haben sich mittlerweile vom reinen Konsumentenmarkt zurückgezogen und bieten heute überwiegend Basistechnologien oder Produkte für bestimmte Fachbereiche an (Kurzweil, Philips etc.).

Erfahrungen im LDS NRW

Im LDS NRW beschäftigt man sich schon sehr lange mit den wichtigsten Produkten zur Spracherkennung. So

wurden bereits frühzeitig Systeme mit diskreter Diktierweise erprobt. Im vergangenen Jahr wurden in umfangreichen Tests die damaligen Produkte der beiden Marktführer IBM und ScanSoft (Dragon) untersucht. Mit beiden konnten nach entsprechender Trainingszeit Erkennungsraten über 95 % erreicht werden. Zur PC-Steuerung und Befehlseingabe erwiesen sich grundsätzlich ebenfalls beide Systeme als geeignet, geringe Unterschiede bestanden dabei in der jeweiligen Bedienung der Produkte.

Die vollständigen Ergebnisse des Vergleichs können von Nutzern der Landesverwaltung auf den Intranetseiten des LDS NRW eingesehen werden (<http://lv.landesintranet.lds.nrw.de/...>).

Zukünftige Entwicklung

Ein besonderer Schwerpunkt der derzeitigen Entwicklung ist das Erkennen „spontaner Sprache“ wie z. B. während eines Interviews. Unter „spontaner Sprache“ versteht man gesprochene Äußerungen, die nicht vor dem Sprechen – wie bei den meisten Diktaten üblich – bereits gedanklich vorformuliert wurden, sondern die spontan beim Sprechen gebildet werden. Hierbei kommt es im Ergebnis oft zu grammatikalisch falschen Sätzen, gefüllt mit Lautäußerungen, Sprechpausen, Korrekturen usw. Derartige Sprechweise ist nur sehr schwer zu analysieren.

Weiterhin soll die Sprecheranpassung, z. B. durch die Bildung von sog. Sprecherprofilen (weiblich, männlich, Jugendlicher, Senior usw.), weiter verbessert werden.

Das größte Problem der Spracherkennung ist jedoch, dass der PC nach wie vor den eigentlichen Sinn der gesprochenen Sätze nicht wirklich erfasst, weshalb die Erkennungsgenauigkeit bei einem Diktat nie volle 100 % erreicht. Kann diese Hürde zukünftig einmal überwunden werden (z. B.

durch neue Algorithmen, Einsatz künstlicher Intelligenz etc.), sind echte Dialoge zwischen Mensch und Computer möglich. Damit würden Sprachanwendungen einen vollkommen neuen Qualitätsstandard erreichen.

Andere Probleme werden jedoch mit hoher Wahrscheinlichkeit bald gelöst werden. So wird die teilweise noch notwendige Trainingsphase, also das „Erlernen“ der Sprachbesonderheiten einer Person bei sprecherabhängigen Systemen, immer weiter auf nur noch wenige Sätze verkürzt werden können. Auch Dialekte und Akzente werden hierbei immer besser verarbeitet. Ein weiterer Schritt auf dem Weg zur umfassenden Sprecherunabhängigkeit.

Ebenso wachsen die mitgelieferten Standardvokabulare weiter an (Stand heute: ca. 1 Mill. Worte inkl. sämtlicher Ableitungen).

Im Bereich der Steuerung und Befehlseingabe wird die Bedeutung von Spracherkennung stark zunehmen. Bereits heute erlaubt nahezu jedes bessere Handy die Sprachwahl. Da gerade in der mobilen Welt (Notebook, Tab-

let-PC, PDA, Palm, Smartphone etc.) die Geräte immer kompakter werden, stößt die Bedienung dieser Geräte über Tastatur und teilweise auch schon über Stift an ihre natürlichen Grenzen.

Studien neuester Entwicklungsprojekte großer Hersteller aus dem Bereich der mobilen Kommunikationstechnik, wie beispielsweise einem PDA in Form einer größeren Armbanduhr oder einem Notebooknachfolger in Form einer Brille, nutzen die Spracheingabe (realisiert auf Chipbasis) als primäre Bedienungsmöglichkeit. Das Problem der persönlichen Authentifizierung (Passwortsatz) ließe sich auf diesem Weg ebenfalls sehr elegant lösen: Das Gerät reagiert – falls gewünscht – nur auf Befehle einer bestimmten Stimme.

Ein weiterer überaus wichtiger Aspekt des Einsatzes von Spracherkennungslösungen fällt unter das Stichwort „Barrierefreiheit“. Die behindertengerechte Bedienung von PCs. Dies betrifft sowohl die Nutzung der genannten Sprachanwendungen als auch die sonstiger Programme (Office-Produkte, Eigenanwendungen) bzw. der Be-

triebssysteme selbst. Auch die Möglichkeit, sich bestehende Texte vom System vorlesen zu lassen (Sprachsynthese, Text-to-Speech) zählt heute zum gewohnten Standard. Das konventionelle Arbeiten mit Tastatur und Maus kann hier mit den heutigen Lösungen vollständig durch Spracheingabe/-steuerung ersetzt werden.

Fazit

Es bleibt festzuhalten, dass die derzeit marktverfügbaren Systeme zur Spracherkennung bereits heute bei entsprechenden Rahmenbedingungen einen sinnvollen und effizienten Einsatz erlauben. In weiten Bereichen wird Spracherkennung unverzichtbar werden. Die ehemalige Fiktion eines PCs, der die Sprache des Menschen „versteh“, wird langsam Realität.



Dieter Bittner
Tel.: 0211 9449-6864
E-Mail: dieter.bittner@lds.nrw.de



Ihre Vorteile

Wir helfen Ihnen, die „richtige“ Archivierung Ihrer Daten zu finden und zu verwirklichen.

Unsere Stärken sind Ihre Vorteile:

- analoge Archivierung seit 35 Jahren
- digitale Archivierung seit 10 Jahren
- selbsttragende Archive auf CD/DVD
- zentrale Archive Online
- Datenverarbeitung und Datenhaltung im Sicherheitsbereich
- Komplettlösungen vom Papier oder von der Datei hin zum digitalen Archiv
- Beratung zu Langzeitstrategien

Wir freuen uns auf Ihre Aufgabenstellung und unterbreiten Ihnen gern unser Angebot.




Archivierung. Eine Dienstleistung des Landesamtes für Datenverarbeitung und Statistik Nordrhein-Westfalen. 2005

Kontakt

Archivierung
Telefon: 0211 9449-2180
E-Mail: rechenzentrum@lds.nrw.de

Weitere Informationen finden Sie unter:
<http://lv.landesintranet.lds.nrw.de/datenverarbeitung/rechenzentrum/archivier/index.html>



© LDS NRW, Düsseldorf, 2005



BR21.001.2005/4

JUDICA/TSJ

Neue Zusammenarbeit mit dem Justizministerium

Seit dem 1. Juni 2005 hat das LDS NRW die Pflege und Weiterentwicklung der IT-Anwendungen JUDICA und TSJ übernommen. Gemeinsam bilden JUDICA und TSJ eine einheitliche Softwareunterstützung für den gesamten Bereich der ordentlichen Gerichtsbarkeit.

Während JUDICA die erforderlichen Daten liefert und verwaltet, kommt TSJ in der Textproduktion zum Einsatz. Aufgrund seiner modularen Struktur und seines hohen Grades an Flexibilität kann JUDICA Funktionen aller Fachbereiche und Gerichtsbarkeiten abbilden und zeitnah, z. B. an Gesetzesänderungen, anpassen.

Die Abkürzung **JUDICA** steht für

Justizunterstützung durch Instanzübergreifende Client-Server Applikation

Justizunterstützung

JUDICA bietet eine umfangreiche Unterstützung für alle Arbeitsschritte, die beim Ablauf eines Gerichtsverfahrens in der jeweiligen Instanz und im jeweiligen Fachbereich auftreten – vom Anlegen eines Verfahrens bis zur Beendigung. Dabei können bereits beendete und schon laufende Verfahren ebenfalls hinzugefügt werden.

Neben der Datenverwaltung der an dem Verfahren beteiligten Personen kommen hierbei der Aktenverwaltung, einschließlich der Führung von Standorthistorien einzelner Aktenbestandteile (Aktenbewegungskontrolle), sowie der Verwaltung der mit einem Verfahren verbundenen Fristen (Wiedervorlagefristen, Fristenkontrolle) eine zentrale Rolle zu.

Diese Funktionalitäten ersetzen zahlreiche vormals von Hand zu führende Listen und Verzeichnisse. Automatisiert werden außerdem die Führung von Zählkarten zur Erfassung statistischer Daten, sowie deren Zusammenfassung und die ausführliche Registerführung, z. B. zur Erstellung einer Liste aller verfahrensbeteiligten Personen.

Darüber hinaus gibt es eine Vielzahl justizfachlicher Funktionen, z. B. zur Berechnung der wirtschaftlichen Grundlage für die Bewilligung von Prozesskostenhilfe oder zur Anrechnung von Untersuchungshaftzeiten auf die Gesamthaftdauer.

Die zentrale Datenhaltung von JUDICA vermeidet auf einfachem Weg die redundante Datenerfassung. Listen mit Daten zentraler Bedeutung (z. B. Behörden, Dolmetscher, Sachverständige und Übersetzer sowie Anwaltslisten) verschaffen den Anwender/-innen einen schnellen Überblick. Umfangreiche Suchfunktionen, die sich benutzerdefiniert einstellen lassen, schaffen hierbei eine wesentliche Erleichterung des Arbeitsablaufs. Immer wieder vorkommende Listen, z. B. Kennziffern von Gerichten und Behörden, werden von JUDICA systemweit vorgegeben. Dadurch werden Unstimmigkeiten vermieden und dank der einheitlichen Darstellung eine effizientere Anwendbarkeit erreicht.

Neben der Verfahrensbearbeitung bietet JUDICA zahlreiche Funktionen für die Gerichtsverwaltung. Hierbei spiegelt JUDICA die interne Struktur des Gerichts und folgt dabei der Aufbauorganisation vom Gericht zur Abteilung zur Organisationseinheit zum einzelnen Bediensteten. Ein wichtiger Aspekt ist die Raumverwaltung. Diese wird von JUDICA bis hin zur Steuerung der elektronischen Anzeige an Sitzungssälen unterstützt.

JUDICA ermöglicht es den Anwender/-innen, einen individuellen, aufgabenbezogenen Terminkalender zu führen. Bei der Anzeige wird zwischen verfahrensabhängigen und verfahrensunabhängigen Terminen unterschieden. Bei der Anzeige von Terminen eines Verfahrens können die anzuzeigenden Detailinformationen nicht durch die Benutzer/-innen geändert werden. Bei der Anzeige von Terminen eines Datums können die Nutzer/-innen die anzuzeigenden Inhalte nach ihren Bedürfnissen auswählen, z. B. alle Termine einer bestimmten Gerichtsabteilung, und die gewählten Einstellungen auch zur Wiederverwendung speichern. Die Suche nach freien Terminen kann sowohl über die Termindauer als auch über den erforderlichen Raum bestimmt werden. Ein gefundener freier Termin lässt sich dann in die Terminerfassung übernehmen.

Über ein Kompetenz- und Rechteverwaltungssystem werden die Zugriffsrechte der Benutzer/-innen auf Daten und Funktionen der Anwendung geregelt. Eine Aufgabe des Administrators besteht in der Betreuung pflegbarer Listen, die den Anwender/-innen während der Verfahrensbearbeitung z. B. mögliche Anreden und mögliche Rollen für Verfahrensbeteiligte zur Auswahl anbieten.

Instanzübergreifend

JUDICA ist nicht auf eine bestimmte Instanz beschränkt, sondern deckt den Bedarf von Amts-, Landes- und Oberlandesgerichten ab, so dass ein Gerichtsverfahren, welches mehrere Instanzen durchläuft, bequem auf elektronischem Weg weitergegeben werden kann. Die Besonderheiten der Instanzen und ihrer Abteilungen werden in verschiedenen Ausprägungen der Software berücksichtigt, die bisher für die Zivil-, Familien-, Insolvenz- und Strafabteilungen realisiert sind. Schnittstellen zur Zusammenarbeit mit weiteren für die Justiz relevanten EDV-Anwendungen, wie etwa zu der Software der Staatsanwaltschaften (MESTa), zum Bundeszentralregister (BZR) und zum Verkehrszentralregister (VZR), sind bereits realisiert.

Client-Server Applikation

Technisch ist JUDICA in einer Dreischicht-Architektur aufgesetzt, bestehend aus der Datenbankschicht, der Applikationsschicht und der Benutzeroberfläche. Die Benutzeroberfläche wurde unter Verwendung der MFC-Klassen unter Visual C++ 6.0 erstellt, die Applikationsschicht in C++ unter Verwendung der Rogue Wave-Klassenbibliotheken. Die Plattform für Oberfläche und Applikation ist Microsoft Windows (NT und 2000). Die Datenbank ist als relationales Datenbank-Management-System realisiert und unabhängig von einem bestimmten Produkt. Als konkrete Ausprägung wird das aktuelle Produkt der Firma Oracle verwendet.

Oberfläche und Applikationsschicht sind zurzeit auf der Client-Seite („Fat Client“) installiert. Dies führt zu einem zu einer verringerten Netzwerkbelastung, zum anderen kann die Anzahl von Serversystemen mit entsprechender Ausfallsicherheit dadurch gering gehalten werden. Die gewählte Programmstruktur bietet jedoch die Mög-

lichkeit, diese Lösung zu einem späteren Zeitpunkt so umzustrukturieren, dass eine Aufteilung der drei Schichten auf drei verschiedene Rechner vorgenommen werden kann. Dabei läuft die Oberfläche zusammen mit einem Rumpf-Client-Programm („Thin Client“) auf dem Arbeitsplatzrechner, die eigentliche Applikation läuft auf einem Windows-Serversystem und die Datenbank bei beiden Lösungen auf einem separaten Datenbankserver. Die Kommunikation zwischen Applikations- und Datenbankschicht wird durch die Verwendung der Rogue Wave C++-Klassenbibliothek ohne den Einsatz von ODBC realisiert.

Die Offenheit und Flexibilität ergibt sich aus der Trennung von Benutzeroberfläche, Applikationsschicht und Datenbank. Die einzelnen Komponenten können anhand der Definition der Schnittstellen unabhängig voneinander entwickelt und angepasst werden.

Textsystem Justiz (TSJ)

TSJ bietet Unterstützung für alle beim Ablauf eines Verfahrens anfallenden Arbeitsschritte bezüglich der Erstellung von Schriftstücken. Es erlaubt, strukturierte Arbeitsanweisungen (z. B. richterliche Verfügungen) zur Anfertigung von Schriftstücken (z. B. den zu einer Verfügung gehörigen Ladungen an beteiligte Personen und entsprechende Rechtsbelehrungen) automatisiert auszuführen. In der Regel wählen die Anwender/-innen hierzu eine Verfügungsvorlage aus einem bestehenden Angebot aus. Die benötigten Detaildaten werden aus JUDICA zur Verfügung gestellt. Andersherum können Daten zu erstellten Schriftstücken in die Schriftstückverwaltung von JUDICA zurück geschrieben werden.

Über die Komponente TSJ-Direkt können auch unmittelbar in JUDICA Schriftstücke erstellt werden, zu denen es keine Verfügungsvorlage gibt, z. B.

wenn ein Termin ohne direkten Zusammenhang zu einer Verfügung eine Einladung erfordert. Das Aussehen und die Struktur der erstellten Reinschriften werden in der Formularverwaltung von TSJ einheitlich festgelegt.

Der Administrator kann in der Ablaufverwaltung Abhängigkeiten zwischen Verfügungen anhand von häufig auftauchenden Konstellationen festlegen, so dass der Anwender bei der Erstellung von Verfügungen und regelmäßig damit verbundenen Folgeverfügungen durch das System geleitet wird.

Die Zusammenarbeit

Die Zusammenarbeit zwischen dem LDS NRW und dem Justizministerium des Landes NRW erfolgt über die Verfahrenspflegestellen der Oberlandesgerichte Köln und Düsseldorf.

Zwischen 2001 und 2005 befand sich JUDICA in der Pilotphase bei den Endanwendern. Aktuell liefern die Verfahrenspflegestellen JUDICA und TSJ bei den Gerichten in Nordrhein-Westfalen aus, so dass in naher Zukunft dem Einsatz in allen 130 Amtsgerichten, 19 Landgerichten und 3 Oberlandesgerichten in Nordrhein-Westfalen an ca. 8 000 Arbeitsplätzen nichts mehr im Wege steht.

Um die komplexen Systeme und Zusammenhänge kennen zu lernen und sicher und schnell auf Anforderungen der Anwender/-innen reagieren zu können, ist eine mehrmonatige Einarbeitungsphase für das LDS-JUDICA/TSJ-Team geplant. Die Verfahrenspflegestellen in Düsseldorf und Köln sammeln Anforderungen und Fehlermeldungen in einer Fehlerdatenbank. Größere Weiterentwicklungen und Änderungswünsche werden in einer gemeinsamen Planungsgruppe besprochen und priorisiert, so dass immer eine klare Aufgabenstellung gegeben ist.

Abb. 1: Baumstruktur des Aktendeckels

JUDICA und TSJ bedeuten für die Gerichte eine wesentliche Arbeitserleichterung und sind nach der Einführung durch den Einsatz in vielen wichtigen Feldern (z. B. Fristenkontrolle oder Aktenverfolgung) aus der täglichen Arbeit der Anwender/-innen nicht mehr wegzudenken. Daher werden besondere Anforderungen an die Stabilität des Systems gestellt. Dies soll durch ein automatisiertes Testverfahren erreicht werden, bei dem die Testfälle gemeinsam vom LDS NRW und den Verfahrenspflegestellen definiert werden.

Da sich die Verfahrenspflegestellen in die Modellierung mit UML (Unified Modeling Language) eingearbeitet haben, wird die Modellierung mit Hilfe von Anwendungsfall-, Sequenz- und auch Klassendiagrammen eine wesentliche Rolle in der Kommunikation zwischen dem LDS NRW und den Kunden spielen.

Dem LDS NRW wächst mit der Weiterentwicklung und Pflege von JUDICA/TSJ eine bedeutende, innovative und dauerhafte Aufgabe zu, die die Zusammenarbeit mit der Justiz weiter stärken wird.

Abb. 2: Aktendeckel – das vertraute Original

Um den Anwenderinnen und Anwendern die Bedienung so intuitiv wie möglich zu gestalten, ist die Oberfläche und Steuerung nahe an andere bekannte Programme angelehnt. Die Baumstruktur des Aktendeckels oder das Sortieren nach bestimmten Spalten in einer Übersicht erinnern an den Windows Explorer, Tabellenkonfigurationen wie das Verschieben von Spalten funktionieren wie in MS Excel®. Das den Anwender/-innen am meisten vertraute Papier, der Aktendeckel, ist dem Original in Aufbau und Funktionalität nachempfunden, bringt aber naturgemäß Vorteile mit, die das Papier nicht bieten kann, z. B. die Verfügbarkeit, Listen ohne Platzbeschränkung, dafür mit Sortier- und Editiermöglichkeiten, keine unleserlichen Handschriften usw.

Während früher erst die entsprechenden Akte vorgelegt werden musste, um eine bestimmte Information in Bezug auf das Verfahren zu erhalten, reicht mit JUDICA oft schon ein Blick auf den (elektronischen) Aktendeckel und ein entsprechender Klick in die Baumstruktur, um Detailinformationen über einen bestimmten Verfahrensbeteiligten, einen Termin oder den Standort einer Akte oder Beiakte zu gewinnen.



Kirsten Rölleke
Tel.: 0211 9449-6907
E-Mail: kirsten.roelleke@lds.nrw.de



Claudia Schröder
Tel.: 0211 9449-6928
E-Mail: claudia.schroeder@lds.nrw.de



Ansprechpartner:
Dr. Peter Gebauer
Tel.: 0211 9449-6902
E-Mail: peter.gebauer@lds.nrw.de

Aufwandsschätzungen in Softwareprojekten

Je mehr der öffentliche Dienst unter finanziellen Druck gerät und Mitarbeiterinnen und Mitarbeiter nicht mehr wie bisher pauschal finanziert für Entwicklungsaufgaben zur Verfügung stehen, desto mehr wird es notwendig, Aufwandsschätzungen bei Entwicklungsprojekten vorzunehmen.

Aufwandsschätzungen sind dabei in unterschiedlichen Zusammenhängen relevant. Sie dienen dazu, anhand konkreter Kostenvoranschläge mit dem Kunden zu verhandeln, Kosten und Nutzen im Projekt zu beurteilen und Fertigstellungstermine zu bestimmen.

Wichtig ist bei solchen Aufwandsschätzungen, dass eine möglichst hohe Schätzgenauigkeit erzielt wird und die Schätzungen in objektiver, nachvollziehbarer Weise erfolgen. Die Schätzungen sollten dabei unabhängig von unrealistischen und unbegründeten Erwartungen sein und auf fundierten Schätzerfahrungen beruhen.

In diesem Artikel werden verschiedene Probleme beim Schätzen aufgezeigt und eine Reihe von Schätzmethoden dargestellt.

Konkrete Umsetzungsmöglichkeiten von Schätzmethoden in der Landesdatenverarbeitungszentrale sollen in einem weiteren Artikel in der nächsten Ausgabe der LDVZ-Nachrichten erörtert werden.

Schätzproblematik

Häufig erweisen sich Aufwandsschätzungen als nicht zutreffend: ca. 25 % aller Softwareprojekte werden vorzeitig abgebrochen, ca. 50 % der „erfolgreich“ abgeschlossenen Projekte haben im Projektverlauf Budget- und Terminüberziehungen hinnehmen müssen.¹⁾ Woher kommt das?

Sicherlich gibt es hierfür eine Vielzahl unterschiedlicher Gründe, und es ist immer der Einzelfall zu berücksichtigen.

Gleichwohl können eine Reihe von Punkten identifiziert werden, die immer wieder zu systematischen Fehlern bei Schätzungen führen. Im Folgenden sollen daher einige Dinge betrachtet werden, die man vermeiden sollte:

1) Quelle: The Standish Group: CHAOS-Report USA, 2001)

Unrealistische und unbegründete Erwartungen der Projektbeteiligten spielen eine entscheidende Rolle. Der Aufwand wird meist unterschätzt, aber so gut wie nie überschätzt.

Häufig überschätzt	Häufig unterschätzt
• Qualifikation des Personals • Produktivität des Personals • positive Auswirkungen von Personalaufstockungen • Produktivität von neuen Methoden und Werkzeugen	• Schulungs- und Einarbeitungsaufwand • Aufwand für Dokumentation • Aufwand für Qualitätssicherung sowie für Fehler-suche und -beseitigung • Kommunikationsprobleme • Aufwand für Projektleitung (Projektkoordination, Projektplanung und -steuerung) • Aufwände im Zusammenhang mit Schnittstellen zu anderen Projekten und Organisationseinheiten

Abb. 1: Überschätzung und Unterschätzung (in Anlehnung an [BundschuhFabry])

Typische Fehler bei Schätzungen, die nachfolgend weiter erläutert werden, sind in Abbildung 2 aufgeführt.

Fünf Arten „nicht zu schätzen“:
• Neue Schätzung = alte Schätzung • Neue Schätzung = alte Schätzung + pessimistische Überziehung • Neue Schätzung = alte Schätzung – bisher aufgetretene Überziehung • Schätzung = erwartete richtige Antwort • Schätzung = erwartete richtige Antwort + X

Abb. 2: Schätzfehler [Barbe]

Diese Schätzfehler können beobachtet werden, wenn ein Projekt bereits einige Zeit läuft und der aktuelle Projektstatus mit der Planung verglichen sowie eine Prognose (Schätzung) für den verbleibenden Projektteil abgegeben werden soll (siehe Beitrag in dieser Ausgabe der LDVZ-Nachrichten „Projektmanagement in der Praxis oder: Wo steht mein Projekt?“ von Thomas Reiff und Jan Mütter).

Häufig wird im ersten Ansatz angenommen, dass die neue Schätzung der alten entspricht, d. h. dass das Projekt „ab jetzt wie geplant“ ablaufen wird. Änderungen an Randbedingungen werden hierbei nicht berücksichtigt. Bei diesem Vorgehen handelt es sich nicht um eine Schätzung, sondern um Raten.

Rechnet man zu den bisherigen Schätzungen einen pessimistischen Wert für eine mögliche Überziehung hinzu, so erhält man zwar eine größere Zahl, die aber wiederum geraten ist, es wird nicht wirklich geschätzt, was noch zu tun ist.

Nicht selten wird in Projekten, die bereits eine Verzögerung eingefahren haben, „geschätzt“, dass der noch verbleibende Aufwand nun um genau jenen Betrag kleiner als geplant ist, der bisher als Verzögerung aufgelaufen ist. Hierbei handelt es sich um eine „blauäugige Schönfärberei“, die vorgaukeln möchte, dass das Projekt zum Schluss „in time“ und „in budget“ enden wird.

Eine ebenfalls verbreitete Methode ist das Einschätzen, welche „Schätzung“ der Auftraggeber gerne hören möchte. Auch hierbei werden nicht die tatsächlich zu erwartenden Aufwände geschätzt. In einigen Fällen wird die „erwartete richtige Antwort“ dann noch um einen gewissen Betrag vergrößert, um anzudeuten, dass die Aufgabe nicht so einfach ist, wie gewünscht. Aber auch hier handelt es sich nicht um eine wirkliche Schätzung.

Den Aspekt der „politischen Schätzungen“ werden wir weiter unten noch ein wenig näher betrachten.

Schätzen ist also offenbar nicht so einfach und im Allgemeinen eine undankbare Aufgabe. Eine erste Empfehlung ist dabei die Berücksichtigung der „Pragmatischen Schätzregeln“ [BundschuhFabry]:

Betont wird hier u. a., dass es wichtig ist, Schätzungen gut zu dokumentieren, um so Erfahrungen im Schätzen systematisch aufzubauen. Je mehr unverfälscht dokumentierte Schätzungen vorliegen, desto genauer können weitere Schätzungen vorgenommen werden.

Genauigkeit der Eingangsinformationen vs. Güte der Schätzung

Die Güte einer Schätzung und damit die Reduzierung von Schätzfehlern ist natürlich in entscheidendem Maße abhängig von dem Wissen über das Schätzobjekt, auf das sich die Schätzung bezieht. Oder anders ausgedrückt: Je genauer die Eingangsinformationen, desto größer die Güte der durchgeführten Schätzung, da mit zunehmender Informationsdichte die Unschärfe beim Schätzen immer weiter abnimmt (siehe Abb. 4).

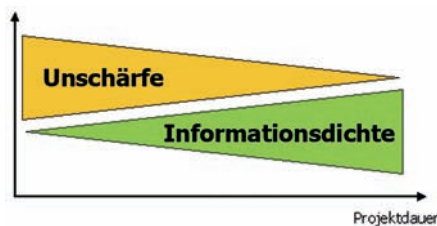


Abb. 4: Zusammenhang zwischen Schätzgüte (Abnahme der Unschärfe) und Informationsdichte

Differenzierung zwischen Aufwandsschätzung und strategischer Angebotserstellung

Oft kommt betriebspolitischen Vorgaben beim Schätzen eine besondere Bedeutung zu. Hier sollte klar getrennt werden zwischen möglichst objektiven Schätzungen einerseits und der Erstellung eines „strategischen Angebots“

andererseits, das aufgrund von Marketingüberlegungen erstellt wird. Hierbei kann es sinnvoll sein, z. B. um neue Märkte zu erschließen, dass eine Leistung zu einem Preis angeboten wird, der unterhalb der Kosten liegt, die dem geschätzten Aufwand entsprechen. Auf der anderen Seite werden natürlich Gewinn bringende Projekte angestrebt.

Eine Aufwandsschätzung kann zwar bestimmte Anhaltspunkte für die Preiskalkulation liefern, letztendlich ist die Preisfestsetzung aber eine Aufgabe der entsprechenden Entscheidungsträger. Wichtig ist, dass stets „vernünftige“ Schätzungen – also unabhängig von irgendwelchen Angebotsaspekten bzw. von „politischen“ Preisen – vorgenommen werden. Insofern können auch Situationen eintreten, in denen angesichts der zu unklaren Schätzgrundlagen überhaupt keine Schätzung möglich ist, sondern lediglich ein „politischer“ Preis auf entsprechender Entscheidungsebene festgelegt werden kann.

Schätzsituation

Bei der Beurteilung der Schätzsituation können Unterscheidungen einerseits hinsichtlich der Projektphase und andererseits bezogen auf die zur Verfügung stehende Informationsdichte gemacht werden.

Pragmatische Regeln (für das Schätzen):

- Je früher Sie schätzen, desto größer ist die Bandbreite der Schätzunschärfe.
- Jede Schätzung ist genauer als keine Schätzung.
- Je besser Sie Ihre Schätzungen dokumentieren, desto größer ist Ihre Chance, Erfahrungen im Schätzen zu erwerben.
- Je mehr dokumentierte Schätzungen zur Verfügung stehen, desto genauer sind Ihre Schätzungen.
- Halten Sie zur Reduktion der Komplexität die Schätzeinheiten möglichst klein und die Arbeitseinheiten möglichst unabhängig.
- Vermeiden Sie den häufigen Fehler, den Kommunikationsfaktor im Team als Aufwandstreiber zu vernachlässigen.
- Denken Sie daran, dass es keine 1:1 übertragbaren Schätzformeln gibt.
- Schätzung ist kein Selbstzweck, sondern dient lediglich der Entscheidungsfindung.

Abb. 3: Pragmatische Schätzregeln [BundschuhFabry]

Projektphase

Soweit das Schätzen vor Projektbeginn – also in einer **Angebotssituation** – erfolgt, wird aufgrund der zu diesem Zeitpunkt vorliegenden Informationen in der Regel nur eine grobe Schätzung möglich sein.

Bei **projektbegleitenden Schätzungen** nimmt zum einen das Wissen über das Schätzobjekt mit der Projektdauer typischerweise immer weiter zu und zum anderen können die Ergebnisse vorangegangener Schätzungen mit einbezogen werden.

Informationsdichte

Wer kennt das nicht: Der Kunde hat – aus welchen Gründen auch immer – nur eine sehr vage Vorstellung davon, was das zu erstellende System leisten soll und dennoch erwartet er, dass ein Auftragnehmer ihm auf den Euro genau beziffern kann, was ihn das Produkt letztendlich kosten wird. Eine Schätzung auf der Grundlage vager Kundenwünsche kann sicherlich auch nur sehr vage sein. Hier sollte keine Scheinobjektivität vorgegaukelt werden. Sinnvoll ist es in solchen Fällen, dass für Aspekte, die noch offen sind, nachvollziehbare Annahmen getroffen und dokumentiert werden und diese dann bei der Schätzung berücksichtigt werden. Für solche Schätzung ist es besonders wichtig, dass nicht ein einzelner Wert als Schätzung angegeben wird, sondern dass eine Bandbreite für unterschiedliche mögliche Projektsituationen beschrieben wird.

Liegen hingegen die Ergebnisse einer detaillierten Systemanalyse vor, so kann die Schätzung auf einer viel genaueren Schätzgrundlage erfolgen und es können auch anspruchsvollere Schätzmethoden verwendet werden.

Schätzmethoden

Heuristische Verfahren

Zunächst betrachten wir die sog. heuristischen Schätzverfahren, also solche Schätzverfahren, die aus der Erfahrung heraus entstanden sind und häufig auf dem im Unternehmen bzw. im Projekt verfügbaren Expertenwissen beruhen und nicht auf fundierten Modellen für die Messung des Systemumfangs. Solche Schätzverfahren werden – in Abgrenzung zu den parametrischen Schätzmethoden – auch als nicht parametrische Schätzmethoden oder auch als Ad-hoc-Schätzverfahren bezeichnet.

Im einfachsten Fall einer **Expertenbefragung** wird die Schätzung durch eine einzelne Person durchgeführt. Dann hängt die Güte der Schätzung natürlich wesentlich von den Erfahrungen und Kenntnissen dieser Person ab. Aber bereits auf dieser Stufe können die Schätzgrundlagen dadurch verbessert werden, dass eine sorgfältige Analyse der Anforderungen und Risiken erfolgt und bewährte Vorgehensweisen angewendet werden. Diese Maßnahmen können dazu beitragen, dass beim Schätzen keine Arbeitspakete übersehen werden.

Eine weitere Verbesserung bringt sicherlich die Einbeziehung mehrerer Experten. Bei einer solchen Experten-schätzung, die vermutlich das am häufigsten angewandte Schätzverfahren darstellt, erfolgt eine Einschätzung des Aufwands durch mehrere fachkundige Personen (**Mehrfachbefragung**). Die Personen sollten aus unterschiedlichen organisatorischen Bereichen kommen bzw. unterschiedliche Aspekte (z. B. aufgrund unterschiedlicher Erfahrungen) in die Schätzung einbringen. Das Ergebnis der Schätzung ergibt sich dann typischerweise als Mittelwert der einzelnen Einschätzungen.

In einer **Schätzklausur** schätzen die Experten gemeinsam in einer Gruppe. Als Unterlagen für die Schätzung die-

nen – neben einer Beschreibung der Projektziele – Projektstrukturpläne, in denen die einzelnen Arbeitspakete definiert sind. Ferner werden – soweit vorhanden – Lasten- bzw. Pflichtenhefte mit hinzugezogen und die Projektrahmenbedingungen (Entwicklungsumgebung, Qualifikation der Mitarbeiter/-innen usw.) berücksichtigt.

Da Schätzungen über alle Arbeitspakete des Projekts hinweg i. d. R. zu zeitaufwändig sind, bietet es sich an, nur wenige repräsentative Referenzkomplexe auszuwählen, deren Arbeitspakete dann einer detaillierten fachlichen Untersuchung und anschließenden genaueren Aufwandsschätzung unterzogen werden. Falls die einzelnen Schätzwerte sehr weit auseinander liegen, werden zwischen den betreffenden Schätzern die Gründe für die abgegebene Schätzung erläutert und diskutiert. Hierbei kann auch geklärt werden, wo u. U. unterschiedliche Annahmen zugrunde gelegt wurden. In gleicher Art und Weise können die einzelnen Arbeitspakete des betrachteten Referenzkomplexes geschätzt werden, wobei jeweils auch der Unsicherheitsgrad der jeweiligen Schätzung beurteilt werden sollte.

Die Übertragung der Schätzergebnisse für den näher untersuchten Referenzkomplex auf die anderen Projektteile erfolgt dann anschließend durch Analogieschluss.

Das Gesamtergebnis der Schätzklausur wird zum Sitzungsende in einem Schätzformular protokolliert.

Ein Vorteil einer solchen Schätzklausur kann darin gesehen werden, dass durch die gemeinsame Diskussion eine von allen Beteiligten akzeptierte Bewertung erzielt wird, die auch besser nach außen vertreten werden kann. Zudem wird das Verständnis für die vermuteten Unsicherheiten geschärft und der Wissensstand im Team angeglichen.

Es muss allerdings darauf geachtet werden, dass die Schätzergebnisse nicht

Delphi-Befragung

1. Projektleiter (bzw. Koordinator) erläutert jedem Experten das anstehende Projekt und übergibt ihm eine Spezifikation und ein Schätzformular.
2. Jeder Experte füllt getrennt das Formular aus (dabei Klärung von Fragen nur mit Projektleiter; keine Diskussion zwischen den Experten).
3. Projektleiter analysiert die Angaben und erstellt Zusammenfassung der Schätzergebnisse sowie Kommentare zu starken Abweichungen der einzelnen Schätzwerte (auf Wiederholformular), Rückmeldung an die Experten erfolgt in anonymisierter Form.
4. Experten überarbeiten ihre Schätzungen unabhängig voneinander.
5. Schritte 3 und 4 werden (ggf. in mehreren Runden) wiederholt bis sich eine hinreichende Annäherung der Schätzwerte ergibt (bzw. der Projektleiter die Ergebnisse akzeptiert).
6. Endgültiges Schätzergebnis = Durchschnittswert der jeweils letzten Überarbeitung der einzelnen Ergebnisse

Abb. 5: Ablauf einer (Standard-)Delphi-Befragung

dadurch verfälscht werden, dass einzelne Personen – insbesondere der Projektleiter – die Schätzklausur dominieren.

Bei der **Delphi-Befragung** werden mehrere Experten schriftlich um ihre Einschätzung des Aufwands gebeten. In zwei oder mehr Befragungsrunden wird ein Schätzergebnis nach und nach angenähert. Es handelt sich also um eine besondere Form der strukturierten Mehrfachbefragung.

Bei der sog. Breitband-Dephi-Methode finden zu Beginn und bei jeder SchätZRunde gemeinsame Sitzungen statt, in denen die Schätzaufgaben bzw. die Zwischenergebnisse der jeweils vorausgegangenen SchätZRunde untereinander diskutiert werden.

Charakteristisch für die Delphi-Methode ist, dass durch die Anonymität dieser Methode²⁾ kein Gruppenzwang zur Konformität besteht und die persönliche Einschätzung eines Experten nicht durch die Dominanz einzelner Personen beeinflusst wird.

Die Delphi-Methode bietet sich besonders für stark innovative Vorhaben an, bei denen keine Ist-Daten aus vergleichbaren Projekten herangezogen werden können.

²⁾ Diese Anonymität ist allerdings bei der Breitband-Delphi-Methode nicht immer vorhanden.

Mit der Delphi-Methode ist allerdings ein nicht gerade geringer Zeitbedarf verbunden, so dass ein Einsatz bei kleineren Projekten nur dann gerechtfertigt ist, wenn die besondere Projektsituation dies erfordert.

In Tabelle 1 findet sich eine zusammenfassende Gegenüberstellung der verschiedenen Varianten der Expertenbefragung.

Eine weitere Verbesserung der Schätzergebnisse kann durch eine **Bandbreitenschätzung** erzielt werden, bei der sowohl von Worst-Case als auch von Best-Case-Schätzungen ausgegangen wird, d. h. es werden zum einen Schätzungen unter pessimistischen Grundannahmen und zum anderen unter optimistischen Grundannahmen durchgeführt. Das Ergebnis ist also nicht ein einzelner Schätzwert, sondern eine Bandbreite zwischen

dem Wert, der bei reibungslosem Projektverlauf erwartet wird (best), und demjenigen, der unter (besonders) ungünstigen Bedingungen zu erwarten ist (worst). Des Weiteren wird der wahrscheinliche Wert (mittel) betrachtet, der unter „realistischen“ Grundannahmen bzw. „normalen“ Bedingungen am ehesten zu erwarten ist.

Derartige Bereichsschätzungen liefern daher ein realistischeres Bild über die Schätzergebnisse, da sie nicht die Scheinobjektivität eines einzigen absoluten Schätzwertes vorgaukeln. Dabei kann die Bereichsschätzung sowohl für die Produktgröße als auch für Aufwand und Zeitdauer eines Projekts durchgeführt werden.

Aus einzelnen Angaben zum Aufwand lässt sich ein Erwartungswert für den Aufwand wie folgt bestimmen:

$$\begin{aligned} \text{Erwartungswert} &= (\text{worst} + \text{best} + 4 \times \text{mittel}) : 6 \\ \text{Standardabweichung} &= (\text{worst} - \text{best}) : 6 \end{aligned}$$

Gleichung 1: Gewichteter Mittelwert und Standardabweichungen³⁾

Der Erwartungswert, der sich aus dem gewichteten Mittelwert der Einzelschätzungen ergibt, ist so definiert, dass mit einer Wahrscheinlichkeit von 50 % der zu erwartende Aufwand unterhalb und mit ebenfalls 50 % Wahrscheinlichkeit oberhalb des Erwartungswertes liegt. Die Angabe nur dieses Werts als Aufwandsschätzung erfordert daher eine recht große Risikoto-

³⁾ Es handelt sich hierbei um den Erwartungswert und die ungefähre Standardabweichung bei einer Betaverteilung.

	Einzelbefragung	Mehrfachbefragung	Delphi-Methode	Schätzklausur
Genauigkeit	ungenau	genau	sehr genau	sehr genau
Aufwand der Schätzung	gering	mittel	groß	sehr groß
Anonymität der Einzelschätzung	–	ja	ja	nein
Mitläufereffekte	–	nein	kaum	ja
Identifikation mit Schätzergebnis	mittel – groß	gering	mittel	groß
Sinnvoller Einsatz	kleine Projekte	mittlere Projekte	große Projekte	große Projekte

Tab. 1: Expertenbefragung zur Aufwandsschätzung [Seibert]

leranz des Auftraggebers. Konservativer ist die Angabe des Erwartungswerts plus Standardabweichung, da der zu erwartende Aufwand mit einer Wahrscheinlichkeit von 84 % diesen Wert unterschreitet. Eine Unterschreitungswahrscheinlichkeit von sogar 98 % wird erreicht, falls der Erwartungswert plus 2-mal Standardabweichung als geschätzter Aufwand angegeben wird.

Da diese Aufwandsbestimmung auf drei Werten beruht, wird sie vielfach auch als Dreipunktmethode bezeichnet⁴⁾.

Bei der Präsentation von Bereichsschätzungen sollten nicht nur die Ergebnisse dargestellt werden, sondern es sollten zusätzlich die jeweiligen Ursachen für mögliche Über- und Unterschreitungen des unter normalen Bedingungen erwarteten Projektaufwands aufgezeigt werden:

Bei einem Projektaufwand von z. B. 300 000 EUR mit einer geschätzten Überschreitung angesichts möglicher ungünstiger Bedingungen von 180 000 EUR und einer geschätzten Unterschreitung bei sehr günstigen Bedingungen von 80 000 EUR könnte dies wie folgt aussehen:

Risiken	Chancen
+ 80 000 EUR, falls sich die Bereitstellung der erforderlichen Basisklassen verzögert	– 40 000 EUR, falls sich die neuen Entwicklungswerkzeuge besser als geplant nutzen lassen
+ 60 000 EUR, falls sich die neuen Entwicklungswerkzeuge nicht wie geplant nutzen lassen	– 40 000 EUR, falls die Einarbeitung der Entwickler deutlich verkürzt werden kann
+ 40 000 EUR, falls die Schnittstellen zum Nachbarsystem XY komplizierter werden	

Bereichsschätzungen bieten also eine realistische Bandbreite für die Projektabwicklung. Die typischerweise stets bei Schätzungen vorhandenen Unsicherheiten werden durch die Angabe von Bereichen verdeutlicht. Dabei können Risiken und auch Chancen aufge-

zeigt und mit in die Schätzung einbezogen werden. Auf diesen Aspekt wird in dem Artikel „Risikomanagement in Softwareprojekten“ von Ulrich Andree und Jan Mütter in dieser Ausgabe der LDVZ-Nachrichten näher eingegangen.

Eine der am häufigsten eingesetzten Schätzmethoden ist eine Variante der **Analogieschätzung**. Dabei werden die zu schätzenden Objekte des neuen Projekts mit bereits bekannten in Verbindung gebracht, was voraussetzt, dass in irgendeiner Weise der Aufwand und die Eigenarten abgeschlossener Projekte gesammelt werden.

Beispiel: Im letzten Projekt haben wir im Schnitt pro Dialogmaske X Personentage benötigt. Jetzt haben wir es mit ca. 75 Dialogmasken zu tun, die im Durchschnitt ungefähr gleich kompliziert sind wie die beim letzten Mal, also schätzen wir den Aufwand für die Dialogmasken auf 75 mal X Personentage.

Bei dieser Methode ist es entscheidend, dass es bereits ein möglichst explizites Erfahrungswissen aus bereits abgeschlossenen Projekten gibt, auf

das bei den Analogieschätzungen zurückgegriffen werden kann. Es hat wenig Sinn, die zu schätzenden Objekte mit den Schätzwerten aus bereits durchgeführten Projekten zu vergleichen, ohne zu überprüfen, inwieweit sich die ursprüngliche Schätzung mit den tatsächlich beobachteten Ist-Werten des abgeschlossenen Vergleichsprojekts deckt.

Beispiel: In einem Projekt soll eine Datenbank mit MySQL entwickelt werden. Zur Schätzung liegen Erfahrungen zur Entwicklung einer Dialoganwendung mit IBM-3270-Masken vor. Was nutzt dem Schätzer diese Information?

Besonders wichtig ist also die Vergleichbarkeit der Schätzobjekte mit Referenzobjekten. Dabei müssen nicht alle Eigenschaften übereinstimmen, vielmehr können auch Unterschiede in Bezug auf bestimmte Merkmale bestehen. Solche Unterschiede müssen dann allerdings explizit berücksichtigt werden, d. h. die Vergangenheitswerte dürfen nicht kritiklos übernommen werden, sondern sind im Hinblick auf die aktuelle Projektsituation in geeigneter Form zu gewichten.

Das Analogieverfahren kann dann besonders sinnvoll verwendet werden, wenn von vielen früheren Entwicklungsprojekten Aufwände bzw. Zeiten und deren Zustandekommen möglichst genau festgehalten wurden und für die Schätzung neuer Projekte zur Verfügung stehen.

Eine weitere Schätzmethode stellt das **Prozentsatzverfahren** dar. Dieses setzt einen genau definierten Entwicklungsprozess voraus, für dessen einzelne Phasen prozentuale Anteilswerte bekannt sind⁵⁾. Bei neuen Projekten kann dann nach Abschluss der ersten Phase der Gesamtaufwand durch Extrapolation des bis hierher angefallenen Aufwands ermittelt werden.

Parametrische Modelle

Die Idee der parametrischen Schätzmodelle ist eine formelbasierte Berech-

4) weitere Bezeichnung hierfür: Pi-mal-Daumen-Methode sowie Beta-Methode

5) Eine solche prozentuale Verteilung der Aufwände auf einzelne Phasen könnte z. B. wie folgt aussehen: 20 % Analyse, 30 % Entwurf, 40 % Realisierung, 10 % Einführung. Beträgt der Aufwand für die durchgeführte Analysephase im betrachteten Projekt z. B. 6 Personenmonate, so ergibt sich dann nach der Prozentsatzmethode ein Gesamtaufwand von 30 Personenmonaten.

nung des zu erwartenden Aufwandes anhand einer Reihe von Eingabeparametern.

In den meisten in der Literatur diskutierten parametrischen Schätzmodellen geht dabei die Systemgröße ganz maßgeblich in die Berechnung des Aufwands ein.

Weitere Parameter, die die Randbedingungen des Systems beschreiben, gehen zudem zum Teil als Faktoren und zum Teil als Exponenten in die Formel ein, so dass sich prinzipiell immer der gleiche Aufbau für die Schätzformel ergibt:

$$\text{Aufwand [Personenmonate]} = \text{Faktoren} \times \text{Systemgröße}^{\text{Exponenten}}$$

Gleichung 2: Allgemeiner Aufbau parametrischer Schätzformeln

Bei den parametrischen Schätzmodellen besteht die Problematik also immer darin, möglichst genau die Systemgröße, die Faktoren und die Exponenten zu bestimmen.

Function-Points

Eine der bekanntesten Methoden zur Bestimmung der Systemgröße ist die FunctionPoint-Methode, die durch die International Function Point Users Group (IFPUG) standardisiert wird.

Kernpunkt ist dabei die Annahme, dass das System ausreichend genau durch die äußeren Schnittstellen (Eingabe (EI), Ausgabe (EO), Abfrage (EQ)) sowie die zu verarbeitenden Daten (interne (ILF) oder externe Datenbestände (EIF)) beschrieben werden kann, und dass die innere Geschäftslogik des Systems für die Abschätzung des Realisierungsaufwands nicht ins Gewicht fällt.

Die Erfahrung vieler Projekte hat dabei gezeigt, dass diese Methode für übliche transaktionsorientierte Anwendungen

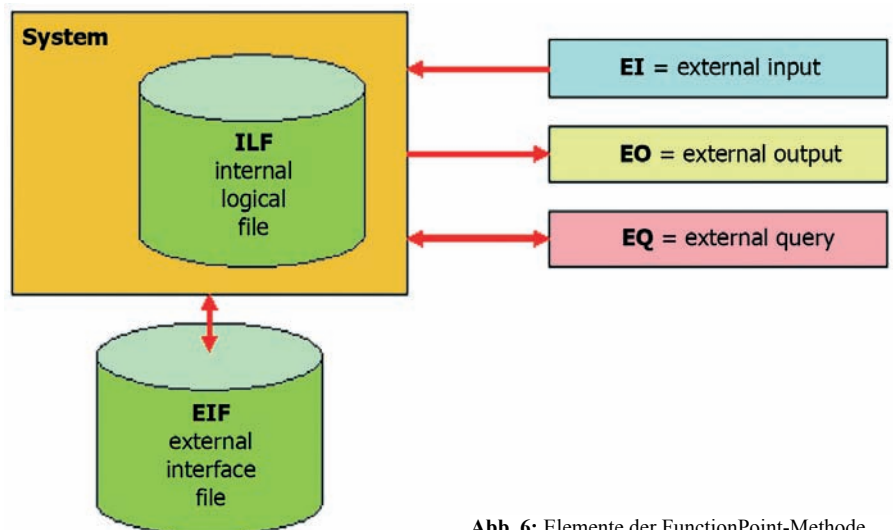


Abb. 6: Elemente der FunctionPoint-Methode

mit einer starken Fokussierung auf die Präsentations- und die Datenhaltungsschicht gut funktioniert, während wissenschaftliche Anwendungen (die im Extremfall eine Zahl als Eingangsparameter erwarten, dann monatelang rechnen, um anschließend eine weitere Zahl auszugeben) mit dieser Methode nicht gut zu beschreiben sind. Details finden sich dazu in vielen Büchern, u. a. in [BundschuhFabry].

Das zu schätzende System wird zunächst „ausgezählt“, d. h. es wird gezählt, wie viele EI's zu bauen sind usw. Außerdem werden die einzelnen Elemente hinsichtlich ihrer Komplexität bewertet, beispielsweise wird nach dem IFPUG-Standard die Komplexität eines external output als „hoch“ bewertet, falls mehr als 20 Attribute mit einer Ausgabe in 2 oder mehr Dateien ausgegeben werden sollen. Die IFPUG stellt zudem ein Bewertungsschema zur Verfügung, nach dem jedem einzelnen Element gemäß seiner individuellen Komplexität eine gewisse Anzahl Funktionspunkte zugeordnet werden.

Aus dem gerade Dargestellten wird deutlich, dass für die Zählung der Funktionspunkte eines Systems bereits ein detailliertes Wissen über das System vorliegen muss, es muss z. B. bekannt sein, wie viele Datenbanktabellen und -attribute das System haben wird. Somit ist diese Methode nicht oder nur sehr eingeschränkt in sehr frü-

hen Projektphasen einsetzbar. Andererseits liefert die Methode in Projektsituationen, in denen sie angewandt werden kann, sehr gute Abschätzungen der Systemgröße, die zudem unabhängig von Randbedingungen wie Entwicklungsstrategie, Technologie u. Ä. sind.

Eine vereinfachte Variante wurde dabei unter dem Namen **Fast Function Points** [Roetzheim] veröffentlicht. Dabei werden nicht mehr alle Details (Anzahl der Attribute in der Datenbank, Eingabefelder auf jeder Dialogmaske, Datenfelder im Ausgabestrom etc.) der zu schätzenden (und meist noch zu konstruierenden) Anwendung gezählt, sondern hier werden erhebliche Vereinfachungen gemacht. Es werden nur noch gezählt:

- Anzahl der **Eingaben** (Datenschnittstellen oder Eingabemasken)
- Anzahl der **Ausgaben**, d. h. Reports entweder als Druckausgaben oder auf dem Bildschirm
- Anzahl der **Datenbanktabellen** (3. Normalform) ohne Information über die Anzahl der Attribute
- Anzahl der **Schnittstellen** (sowohl als Files, Datenbanken, APIs oder sonstiges). Dabei wird genau genommen die Anzahl der „Satzarten“ gezählt und nicht die Anzahl der technisch implementierten Schnittstellen
- Anzahl der **Systemnachrichten** (synchrone oder nahezu synchrone Transaktionsmeldungen in das oder aus dem System)

Diese weniger detaillierten Informationen liegen meist auch schon in einem frühen Entwurfstadium mit einer ausreichenden Stabilität vor, so dass sie die Größe des Systems adäquat beschreiben können.

Untersuchungen an durchgeführten Projekten haben gezeigt, dass die Abweichung der Werte der Fast-Function-Point-Methode von denen der vollen FunctionPoint-Zählung nach [IFPUG] kleiner als 10 % ist.

Weitere Varianten der FunctionPoint-Methode sind z. B. ObjectPoints, DataPoints oder InternetPoints, die jeweils speziell für objektorientierte, datenbankzentrierte oder Webprojekte adaptiert sind.

Von der Systemgröße zur Aufwandsschätzung

Im nächsten Schritt der parametrischen Schätzungen muss nun die Systemgröße mit Hilfe der oben dargestellten Formel mit dem zu erwartenden Aufwand in Verbindung gesetzt werden. Dabei berücksichtigen die Faktoren und Exponenten jene Aspekte, die über die Systemgröße hinausgehend bei einer Aufwandsschätzung zu berücksichtigen sind.

Bereits in den 1980er-Jahren wurde das **Constructive Cost Model (COCOMO)** von Barry Boehm entwickelt [Boehm], das später zum Modell COCOMO II erweitert wurde.

Dieses Modell legt die oben dargestellte Formel zugrunde und beschäftigt sich insbesondere mit der Kalibrierung der Faktoren und der Exponenten.

Die Aufwandsschätzung ist recht empfindlich gegenüber Veränderungen der im COCOMO II-Modell berücksichtigten Faktoren (siehe Abb. 7). Diese können so verändert werden, dass der Gesamt-Multiplikator zwischen 0,05 und 115 variiert.

Faktoren	
Produkt • erforderliche Zuverlässigkeit • Datenbankgröße • Produkt-/Modulkomplexität • Wiederverwendbarkeit • Dokumentationsumfang	Personal • Systemanalysefähigkeiten • Programmierfähigkeiten • Anwendungserfahrung • Plattformerfahrung • Sprach- und Toolererfahrung • personelle Kontinuität
Plattform • Rechnerzeitnutzung • Hauptspeichernutzung • Plattform-Änderungsdynamik	Projekt • Nutzung von Werkzeugen • standortübergreifende Teamarbeit • verfügbare Projektdauer

Abb. 7: Aufwandsmultiplikatoren (Quelle: [Seibert])

Das COCOMO-Modell gibt dem Schätzer dazu detaillierte Hilfen an die Hand, um diese Faktoren richtig zu justieren.

Die Kostentreiber „Erfahrung im Produktbereich“, „Entwicklungsflexibilität“, „Ausgereiftheit bzw. Risikofreiheit des Entwurfs“, „Zusammenarbeit zwischen den Projektbeteiligten“ sowie der „Reifegrad des Softwareentwicklungsprozesses“ gehen in COCOMO II als Exponenten in die Formel ein und können zwischen 0,91 und 1,23 verändert werden.

und andererseits, dass die Faktoren und Exponenten akkurat eingestellt werden.

Die Güte der Schätzungen kann zudem nochmals deutlich verbessert werden, wenn die von COCOMO zunächst mitgelieferten generischen Modellparameter mit Hilfe eigener Ist-Daten aus bereits durchgeführten Projekten kalibriert werden.

Inzwischen wurden eine Reihe von Varianten und Anpassungen von COCOMO II auf spezielle Anwendungsgebiete veröffentlicht:

Modell	Anwendungsgebiete
COCOMO II	Das ursprüngliche CONstructive COSt MODEL.
Agile COCOMO	Eine Adaption von COCOMO für agile Projekte. Dabei werden insbesondere analogiebasierte Schätzmethoden eingesetzt.
COCOTS	Constructive COTS (Custom Off The Shelf) beschäftigt sich mit der Aufwands-, Kosten- und Terminschätzung von Projekten, bei denen gekaufte Fremdkomponenten verwendet und integriert werden sollen.
COQUALIMO	Ein Schätzmodell, das sich mit der Abhängigkeit von Kosten, Termin und Qualität beschäftigt.
CORADMO	Dieses Modell ist auf Rapid-Application-Development-Projekte abgestimmt.
COPROMO	Das CONstructive PROductivity improvement MODEL konzentriert sich auf die Kostenschätzung von Produktivitätssteigerungen.
COPSEMO	Constructive Phased Schedule & Effort Model
COSYSMO	CONstructive SYStems Engineering Cost Model

Abb. 8: Varianten [COCOMO]

COCOMO II ist ein in vielen tausend Projekten erprobtes Modell zur Aufwandsschätzung. Entscheidend für die Güte der Schätzung ist hierbei einerseits die Genauigkeit, mit der die Größe bzw. Komplexität des Systems bekannt ist,

Es gibt zur Unterstützung von Aufwandsschätzungen eine Reihe von nützlichen Werkzeugen, auf die an dieser Stelle nicht näher eingegangen werden soll (als Anregung siehe z. B. [CostXpert] und [COCOMO]).

Aufwandsschätzungen in agilen und evolutionären Projekten

Eine der häufigsten Ursachen für Projektabbrüche sowie Termin- und Kostenüberschreitungen stellen immer wieder die unklar definierten bzw. häufig wechselnden Anforderungen dar. Wie bereits im Beitrag „Agil und extrem (2) – praktische Erfahrungen mit agilen Methoden“ in den LDVZ-Nachrichten 1/2005 aufgezeigt, wird man jedoch in der Vielzahl der Projekte mit sich häufig ändernden Anforderungen leben müssen, d. h. es ist eine Illusion anzunehmen, dass man zu Beginn eines Projekts die Anforderungen einmal festlegen kann und diese dann bis zur Realisierung des Systems stabil bleiben. Insofern stellt sich dann eher die Frage, wie man möglichst geschickt mit solchen sich ändernden Anforderungen umgehen kann. Agile und evolutionäre Vorgehensweisen haben sich in derartigen Situationen inzwischen vielfach bewährt.

Die Aufwandsschätzung in agilen Projekten läuft dann typischerweise wie folgt ab:

- Zu Projektbeginn wird lediglich ein grober Rahmen für das zu entwickelnde System abgesteckt, der allen Beteiligten als grobe Orientierung dient.
- Dieser bildet die Grundlage für den Budgetrahmen.
- Der Funktionsumfang wird offen gehalten.
- Die Realisierung der vom Kunden gewünschten Funktionalität erfolgt nach dem Prinzip des Time-Boxings (d. h. es gibt feste Iterationszyklen).
- Der Kunde legt fest, welche Funktionen in der nächsten Iteration realisiert werden sollen und zwar auf Basis der von den Entwicklern abgeschätzten Realisierungsaufwände für die einzelnen Funktionen.
- Dabei besteht der Anspruch, in jeder Iteration das optimal Mögliche für den Kunden umzusetzen.

Wenn allerdings neue Anforderungen hinzukommen, ohne dass dafür vorhandene Anforderungen entfallen oder reduziert werden, werden sich dadurch auch Budget und Termin ändern. Also muss eine neue Aufwandsschätzung sowie eine möglichst akkurate Positionsbestimmung durchgeführt werden (siehe auch Artikel „Projektmanagement in der Praxis oder: Wo steht mein Projekt?“ von Thomas Reiff und Jan Mütter in dieser Ausgabe der LDVZ-Nachrichten).

Akzeptiert der Kunde das nicht, sondern möchte er den Aufwand und damit die Kosten weiterhin zum Projektbeginn auf Euro genau beziffert haben, dann bekommt er auch genau das, was er zum Beginn des Projekts bestellt hat. Eine Berücksichtigung von Änderungswünschen oder neuen Erkenntnissen im Projektverlauf wird dadurch sehr behindert, und das umso mehr, je größer das Projekt ist.

Zusammenfassung

In diesem Artikel wurden verschiedene Schätzmethoden aufgezeigt. Hieran kann man erkennen, dass inzwischen eine ausgereifte Methodik der Aufwandsschätzung existiert. Bis derartige Methoden allerdings ihren Einzug in die Projektpraxis finden werden, wird noch einiges an „Überzeugungsarbeit“ (und auch an Anpassung der Methoden an die jeweiligen Projektsituationen) notwendig sein. Dem systematischen Aufbau von Erfahrungswissen über eigene frühere Projekte (Erfahrungsdatensammlungen) sowie der Unterstützung der einzelnen Entwicklungsbereiche bei der Durchführung von Schätzungen kommt hierbei eine wichtige Rolle zu.

Wir möchten die Leserinnen und Leser anregen, uns ihre Einschätzungen zu den hier vorgestellten Schätzproblematiken und Schätzmethoden mitzuteilen und uns von eigenen Erfahrungen zu berichten.

Literatur/Links

- [AndreeMütter] Ulrich Andree, Jan Mütter, „Risikomanagement in Softwareprojekten“, LDVZ-Nachrichten, 2/2005
- [Barbe] Benny Barbe, „Planungssicherheit für IT-Projekte“, Seminar CostXPert, 2005
- [Boehm] Barry Boehm et.al. „Software Cost Estimation with COCOMO II“, Prentice Hall, 2000
- [Bundschuh] Manfred Bundschuh, „Einsatz und Nutzen der Function-Point-Methode“, in Projektmanagement Aktuell, 1/2005
- [BundschuhFabry] Manfred Bundschuh, Axel Fabry „Aufwandsschätzung von IT-Projekten“, mitp Verlag, 2004
- [COCOMO] COCOMO-Homepage an der University of Southern California <http://sunset.usc.edu/research/cocomosuite>
- [CostXpert] Cost Xpert Group, Inc. <http://www.costxpert.com>
- [IFPUG] International Function Point Users Group <http://www.ifpug.org>
- [MüttervonHagen] Jan Mütter, Ulrich von Hagen „Agil und extrem (2) – Praktische Erfahrungen mit agilen Methoden“, LDVZ-Nachrichten, 01/2005
- [ReiffMütter] Thomas Reiff, Jan Mütter, „Projektmanagement in der Praxis oder: Wo steht mein Projekt?“, LDVZ-Nachrichten, 2/2005
- [Roetzheim] William H. Roetzheim, „Introduction to Fast Function Points“
- [Seibert] Prof. Dr. Siegfried Seibert, „Aufwandsschätzung von Software-Projekten“, Seminar TAE Esslingen, 2001



Dr. Jan Mütter
Tel.: 0211 9449-5434
E-Mail: jan.muetter@lds.nrw.de



Dipl.-Inf. Ulrich von Hagen
Tel.: 0211 9449-6706
E-Mail: ulrich.von-hagen@lds.nrw.de

Risikomanagement in Softwareprojekten

Das Projekt beginnt. Im Projektplan ist der Weg zum Ziel dargestellt. Jetzt braucht das Team nur noch die geplanten Aktivitäten durchzuführen und schon ist das Projektziel kurze Zeit später erreicht.

Kennen Sie solche Projekte? Wir auch nicht.

Die Realität ist immer anders als geplant und manchmal kommt es derart anders als erwartet, dass das Projekt in große Schwierigkeiten gerät.

Kann man sich überhaupt vor Risiken schützen? Wenn ja, wie geht das und was sind überhaupt Risiken? Und muss ich wirklich Risikomanagement betreiben? Das kostet doch nur Aufwand.

Nach Ansicht der Autoren de Marco und Lister gilt:

„Risikomanagement ist Projektmanagement für Erwachsene.“
[deMarcoLister]

Warum wir das ähnlich sehen und was Risikomanagement bedeutet, soll in diesem Beitrag erläutert werden.

Risiken und Chancen

Risiken und Chancen eines Projektes sind enge Verwandte. Dies macht schon die folgende von der Gesellschaft für Projektmanagement verwendete Definition deutlich [PMF]:

„**Projektrisiken** sind mögliche Ereignisse oder Situationen mit negativen Auswirkungen (Schäden) auf das Projektergebnis insgesamt, auf beliebige einzelne Planungsgrößen oder Ereignisse, die neue unvorhergesehene und schädliche Aspekte aufwerfen können.“

Eine **Chance** ist demnach ein mögliches Ereignis oder eine Situation mit positiven Auswirkungen auf das Projektergebnis insgesamt oder auf Teile davon.

Somit sind Chancen und Risiken zwei Seiten ein und derselben Medaille.

Nun sind wesentliche Charakteristika von Projekten die „Einmaligkeit“ bzw. „Neuartigkeit“ sowie eine gewisse „Komplexität“ (siehe DIN 69901), welche ein Projekt vom „Tagesgeschäft“ unterscheiden.

Natürlicherweise erhofft man sich von der Durchführung eines Projektes einen gewissen Vorteil, also eine Chance. Demgegenüber stehen aber immer die zwangsläufig mit einem Projekt verbundenen Risiken.

Patzak und Rattay haben diese Verknüpfung zwischen Risiken und Chancen in dem folgenden Diagramm ausgedrückt [PatzakRattay]:

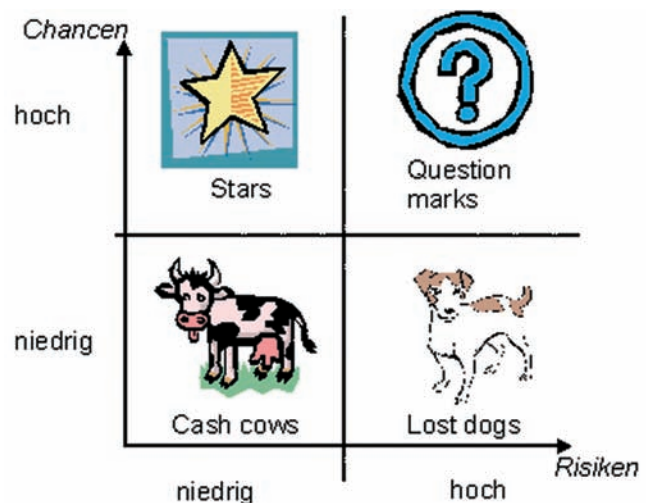


Abb. 1: Zusammenhang zwischen Risiken und Chancen [PatzakRattay]

Dabei sind die **cash cows** jene Projekte bzw. Aufgaben, die ohne großes Risiko ihre Investitionen wieder hereinholen. Von Projekten mit hohen Risiken, aber ohne echte Chancen, den sog. **lost dogs** sollte man die Finger lassen. Die echten **Stars** unter den Projekten, zumindest aus Vertriebsicht, sind jene mit hohen (Gewinn-)Chancen, aber nur geringem Risiko. Aus Innovationssicht sind es aber die **question marks**, die interessant sind.

„Wenn ein Projekt kein Risiko birgt ... lassen Sie die Finger davon.“
[deMarcoLister]

Innovative Projekte, die z. B. in neue Technologien vordringen oder neue Branchen zum Ziel haben, bergen natürlich wesentlich höhere Risiken, als Projekte, die das Unternehmen bereits seit vielen Jahren durchführt. Andererseits ergeben sich hierdurch auch neue Chancen, z. B. in neue Märkte vorzudringen oder neue Kunden zu gewinnen.

Hier muss man sich im Einzelfall sehr genau überlegen, ob die Chancen, die sich das Unternehmen von dem Projekt erhofft, die zu erwartenden Risiken überwiegen oder nicht.

Der Risikomanagement-Prozess

Der Prozess des Risikomanagements vollzieht sich in mehreren Stufen. Dabei wird dieser Prozess nicht nur einmal durchlaufen, sondern im Lebenszyklus eines Projektes viele Male. Ein Projekt zeichnet sich gerade dadurch aus, dass in den verschiedenen Stufen des Lebenszyklus auch verschiedene Risiken (und Chancen) liegen. Beispielsweise wird ein Projekt in der Phase, in der Anforderungen erhoben werden, von ganz anderen Risiken bedroht (z. B. Vollständigkeit, Stabilität und Konsistenz der Anforderungen) als in der Phase, in der das Produkt in die Produktion übergeleitet wird (z. B. Rollout der Infrastruktur).

In den folgenden Abschnitten sollen nun die einzelnen Schritte des Risikomanagement-Prozesses näher betrachtet (siehe auch [Gernert]) und mit Anwendungserfahrungen aus realen Projekten der Landesdatenverarbeitungszentrale (LDVZ) erläutert werden.

1 Identifizieren von Risiken

Der erste Schritt des Risikomanagements ist das Identifizieren von Risiken. Nur die Risiken, die man auch kennt, kann man managen. Daher ist es wichtig, bereits in einer sehr frühen

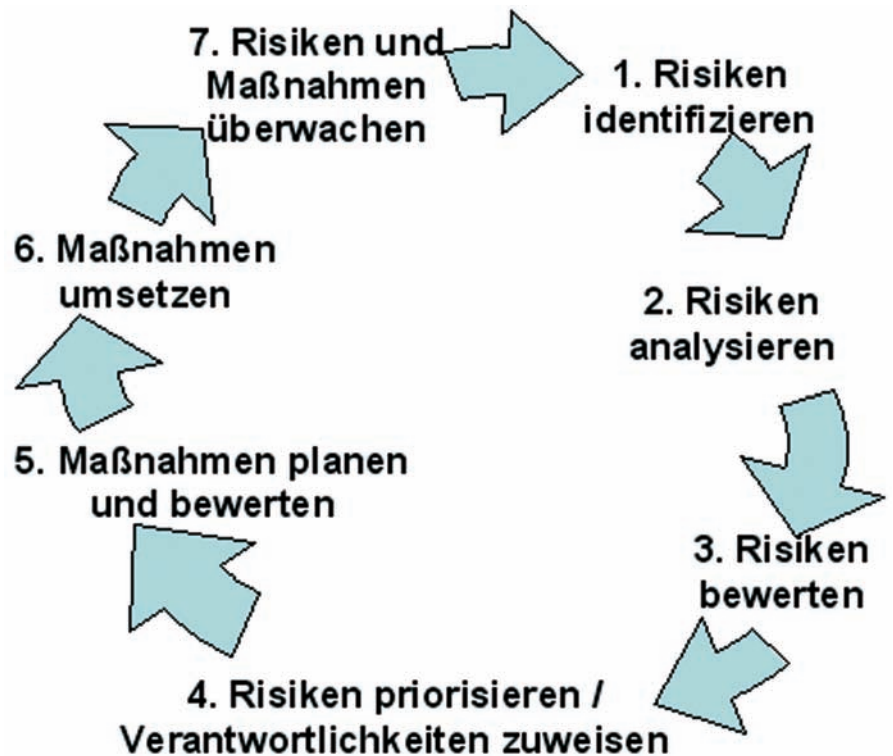


Abb. 2: Der Risikomanagement-Prozess

Phase des Projektes die Risiken zu identifizieren.

Risiken zu identifizieren ist für einen Projektleiter unbequem, muss er sich doch mit Themen auseinandersetzen, die den Projektplan gehörig durcheinander bringen (können). Andererseits ist es fahrlässig und unverantwortlich, die Augen vor möglichen Risiken verschließen zu wollen. Der Eintritt eines nicht erkannten Risikos, auf das das Projekt dann ja auch nicht vorbereitet ist, kann einen erheblichen Schaden für das Projekt und das gesamte Unternehmen bedeuten.

Eine wichtige Methode für die Risikoidentifikation ist die Brainstorming-Sitzung. An dieser Sitzung sollten Projektbeteiligte aus allen Bereichen beteiligt werden, d. h. Projektmitglieder aus den unterschiedlichen Teilprojekten, Stakeholder, der Auftraggeber, das Management. Wichtig ist, dass alle Projektbeteiligten die Möglichkeit haben, sich einzubringen.

In der Sitzung werden jetzt alle Risiken, die den Beteiligten einfallen, ge-

sammelt. Es sollte noch keine Diskussion über die Wichtigkeit der Risiken geführt werden, sondern zunächst eine möglichst vollständige Sammlung aller Bedrohungen entstehen.

Die Risikoidentifikation sollte in regelmäßigen Abständen wiederholt projektbegleitend durchgeführt werden. In einem Risikoworkshop am Anfang des Projektes können nicht alle Risiken für die gesamte Projektlaufzeit erkannt werden.

Tom DeMarco und Timothy Lister [DeMarcoLister] schlagen dabei das „Katastrophen-Brainstorming“ vor. Dabei werden zuerst nicht die Risiken betrachtet, sondern es werden alle möglichen katastrophalen Szenarios, in die das Projekt geraten kann, gesammelt. Dann werden die Katastrophenszenarien analysiert und die dahinterliegenden Ursachen aufgedeckt, die als Risiken das Projekt bedrohen.

Eine weitere Methode ist die Nutzung von Risikolisten. Die einschlägige Literatur enthält umfangreiche Listen mit Risiken, die ein (Software-)Pro-

jekt bedrohen. Dabei kann die Nutzung dieser Listen aber nicht den Brainstorming-Workshop ersetzen, da jedes Projekt durch seine Einmaligkeit auch durch einmalige Risiken und durch einmalige Chancen gekennzeichnet ist, die man nur durch vorgefertigte Risikolisten nicht erkennen kann. Diese Listen sind aber ein gutes Hilfsmittel, um die eigene Risikoliste zu überprüfen.

Das Ergebnis der Risiko-Identifizierung ist eine Liste mit allen möglichen Risiken, die das Projekt bedrohen. Da in der Regel sehr viele Risiken gefunden werden, sollte diese Liste geordnet werden, beispielsweise nach Projektphasen oder anderen sinnvollen Klassifizierungsmerkmalen.

Mit dieser Liste kann der zweite Schritt, die Analyse in Angriff genommen werden.

2 Risiken analysieren

Im zweiten Schritt müssen die gefundenen Risiken detailliert analysiert werden. Alle Risiken haben unterschiedliche Eintrittswahrscheinlichkeiten und bei Eintritt unterschiedliche negative Auswirkungen. Dadurch ergibt sich, dass die Risiken unterschiedlich behandelt werden müssen. Im Analyse-Schritt müssen die wichtigsten, d. h. die bedrohlichsten, Risiken gefunden werden, um insbesondere für diese Risiken Maßnahmen zu finden und Verantwortlichkeiten zuzuweisen.

Bei jedem Risiko werden entsprechend die Eintrittswahrscheinlichkeit und die Tragweite geschätzt. Unter Tragweite versteht man dabei die negativen Auswirkungen eines eingetretenen Risikos, d. h. die Auswirkungen auf die Kosten, auf die Qualität oder den Leistungsumfang des Produktes und auf das Erreichen des Fertigstellungstermins. Die Abschätzung dieser Werte ist natürlich mit einer Unsicherheit verbunden. Sie

1. Ausschnitt aus der Risikoliste eines LDVZ-Projekts			
ID	Kurzbeschreibung	Beschreibung	Risikoklasse
005	Änderungen an bereits fertiggestellten Komponenten	Durch Änderungen an bereits fertiggestellten Komponenten werden ungeplante Überarbeitungsaufwände erforderlich, die das Budget und den Liefertermin gefährden.	Anforderungen
018	Schlüsselpersonen fallen aus	Der längerfristige Ausfall von Schlüsselpersonen im Projekt führt zu Verzögerungen und gefährdet damit Budget und Liefertermin	Personal
042	Gesetzesänderungen	Durch Gesetzesänderungen werden ungeplante Aufwände erforderlich, die das Budget und den Liefertermin gefährden.	Anforderungen
098	Serverausfall	Durch den Ausfall des Servers geht der Zugriff auf das Entwicklungs-Repository verloren. Dies führt zu Verzögerungen und damit zu einer Gefährdung von Budget und Liefertermin.	Technik

sollte daher, genau wie die Risiko-Identifizierung in einem größeren Kreis von Beteiligten vorgenommen werden.

Andererseits gilt auch hier, dass eine Schätzung immer noch besser ist als keine Schätzung.

Die Eintrittswahrscheinlichkeit kann beispielsweise in Stufen (z. B. niedrig – mittel – hoch) oder in Prozent geschätzt werden. Die Schätzung in Prozent erlaubt es, die Risiken feiner abzustufen, täuscht andererseits aber auch eine Genauigkeit vor, die es so in der Regel nicht gibt.

Die Tragweite für die unterschiedlichen Bereiche (siehe oben) kann in Stufen geschätzt werden (z. B. niedrig – mittel – hoch). Diese Stufen sollten vor der Schätzung genau definiert werden, so dass alle Beteiligten auf derselben Skala schätzen. Diese konkreten Definitionen sind sehr stark projektabhängig und können daher nicht generell vorgegeben werden. Bei unterschiedlichen Projekten ist beispielsweise das Risiko eines Verlusts von 100 000 Euro je nach Projektgröße unterschiedlichen Tragweiten zuzuordnen.

In einem Projekt mit einem Budget von 100 000 Euro kommt der Eintritt eines solchen Risikos einem Totalverlust gleich, während die Erprobung einer neuen Technologie, die eine Investition von 100 000 Euro erfordert, in einem großen Projekt möglicherweise vorteilhaft ist, wenn dadurch z. B. die Chance besteht, Entwicklungsaufwände zu reduzieren oder das Produkt leistungsfähiger zu machen.

3 Bewertung von Risiken

Ziel der Bewertung ist es, aufgrund der Risikoanalyse eine Priorisierung der Risiken vornehmen zu können, um die wichtigsten Risiken zu erkennen. Dafür wird aus den bei der Analyse ermittelten Kennzahlen der so genannte Risikowert gebildet. Dieser Risikowert lässt sich nach verschiedenen Methoden errechnen und erlaubt eine direkte Vergleichbarkeit der Prioritäten von Risiken.

Eine Möglichkeit der Errechnung des Risikowerts ist die im Projekt OBELIX eingesetzte Punktemethode [Harrant Hemmrich]. Dabei wird die Eintrittswahrscheinlichkeit in Prozent ge-

schätzt, die Auswirkung auf den Termin und die Auswirkung auf das Produkt werden auf einer diskreten Skala von niedrig = 1, mittel = 5 bis hoch = 10 angegeben.

Der Risikowert errechnet sich nach diesem Verfahren folgendermaßen:

Risikowert
 = (Eintrittswahrscheinlichkeit * 10)
 + Auswirkung auf den Termin
 + Auswirkung auf das Produkt.

Dabei unterscheiden Harrant und Hemmrich noch die Auswirkungen auf den Termin und auf die Kosten. Diese Unterscheidung wird im Projekt OBELIX nicht gemacht, da die Kosten fast ausschließlich durch den Personaleinsatz entstehen und so die Gesamtkosten fast identisch mit den Personalkosten sind.

Das Ergebnis ist in unserem Fall für jedes Risiko ein Wert zwischen 1 (= unbedeutend) und 30 (= extremes Risiko).

Eine andere Möglichkeit die Auswirkungen zu schätzen, ist die Multiplikationsmethode [Thaller]. Dabei wird die Eintrittswahrscheinlichkeit in Prozent und die Schadenshöhe in Euro bei Eintreten des Risikos geschätzt. Der Risikowert läßt sich nach dieser Methode folgendermaßen errechnen:

Risikowert
 = Schadenshöhe * Eintrittswahrscheinlichkeit

Dieser Risikowert kann so interpretiert werden, dass Rückstellungen in dieser Höhe im Projektbudget berücksichtigt werden sollten.

2. Ausschnitt der Risikobewertung eines LDVZ-Projekts

ID	Kurzbeschreibung	Eintrittswahrscheinlichkeit	Auswirkung auf Termin	Auswirkung auf Produkt	Risikowert
005	Änderungen an bereits fertiggestellten Komponenten	0,75	mittel	mittel	17,5
018	Schlüsselpersonen fallen aus	0,5	hoch	mittel	20,0
042	Gesetzesänderungen	0,25	mittel	mittel	12,5
098	Serverausfall	0,25	niedrig	niedrig	4,5

3. Kategorie und Risikowerte [HarrantHemmrich]

Kategorie	Beschreibung	Wert
Niedrig	Es werden keine speziellen Maßnahmen zur Risikovermeidung getroffen, die Risiken werden von dem entsprechenden Arbeitspaketverantwortlichen weiterverfolgt.	0 – 9
Mittel	Es werden Maßnahmen zur Risikovermeidung und -minderung umgesetzt. Das Controlling findet regelmäßig auf Projektstatusbesprechungen statt.	10 – 24
Hoch	Dies sind die Kernrisiken des Projekts. Es müssen zwingend Maßnahmen zur Risikovermeidung und -minderung umgesetzt werden. Zusätzlich zum Controlling in den Projektstatusbesprechungen müssen diese Risiken im Lenkungsausschuss als dem höchsten Gremium des Projekts überwacht werden.	25 – 30

4 Risiken priorisieren und Verantwortlichkeiten zuweisen

Anhand des ermittelten Risikowerts lassen sich die Risiken jetzt priorisieren und Verantwortlichkeiten zuweisen. Eine Möglichkeit zur Priorisierung wird dafür in Übersicht 3 gezeigt.

Wichtig ist außerdem, bei allen Risiken die Verantwortlichkeiten zu untersuchen. Man muss feststellen, ob die Verantwortlichkeiten für das Risiko, bzw. für Gegenmaßnahmen im Projekt selber liegen (internes Risiko) oder außerhalb des Projektes (externes Risiko). Liegt die Verantwortung im Projekt selber, können Verantwortliche für die Maßnahmen direkt benannt werden. Liegt die Verantwortung außerhalb, muss das Risiko an eine übergeordnete Entscheidungsinstanz eskaliert werden, die dann über den weiteren Umgang mit dem Risiko entscheidet.

5 Maßnahmen planen und bewerten

In diesem Schritt werden Risikostrategien für den Umgang mit den gefundenen Risiken, insbesondere den Risiken mit hohem Risikowert, geplant. Wichtig ist dabei, dass es für die meisten Risiken mehrere Strategien gibt, diese aber unterschiedlich wirksam sind und unterschiedlich viel kosten. Diese Maßnahmen können sich gegenseitig beeinflussen oder sogar beeinträchtigen. Da die Maßnahmen häufig auch Auswirkungen auf den Projektplan haben, müssen sie gegebenenfalls mit dem Lenkungsausschuss des Projekts abgestimmt werden.

Harrant und Hemmrich [Harrant Hemmrich] unterscheiden dabei vier verschiedene Risikostrategien:

Risiko vermeiden

Ein Risiko zu vermeiden bedeutet, das Risiko komplett zu beseitigen oder aber

das Projekt vor den Auswirkungen zu schützen. Grundlage dafür ist, dass der eigentliche Auslöser des Risikos bekannt ist. Ein Risiko zu vermeiden ist die beste, aber meist auch die teuerste Strategie.

Beispiel:

Das Risiko einer Client-Server-Anwendung, bei der es zu Problemen bei der Softwareverteilung auf die Clients kommt, lässt sich vermeiden, wenn statt der CS-Lösung eine Web-Anwendung erstellt wird, die vollständig auf einem Server läuft.

Kosten: Sicherlich war die Entscheidung für eine CS-Lösung wohlüberlegt. Dann muss sehr genau geprüft werden, welche Kosten bei der Umstellung auf eine Web-Lösung anfallen, und ob diese Kosten den eigentlichen Risikowert nicht bei weitem übersteigen.

Risiko transferieren

Das Risiko bzw. die Verantwortung für Schäden bei eingetretenem Risiko werden dabei auf Andere verlagert (z. B. indem man eine Versicherung abschließt). Diese Maßnahme schützt allerdings nicht davor, dass das Risiko trotzdem eintritt.

Beispiel:

Das Risiko in einem Software-Projekt, dass der zentrale Repository-Server ausfällt und somit die Arbeit der Entwickler unterbrochen wird, wird auf das Rechenzentrums-Team transferiert, da dieses für den Serverbetrieb zuständig ist.

Kosten: Das Rechenzentrums-Team wird dem Entwickler-Team die Bereitstellung und Administration in Rechnung stellen. In der Regel dürften diese Kosten weit unterhalb des möglichen Schadens liegen.

Risiko vermindern

Bei dieser Strategie versucht man, entweder die Eintrittswahrscheinlichkeit eines Risikos oder aber die Schadenshöhe zu vermindern. Die Verminderung der Eintrittswahrscheinlichkeit ist dabei die bessere, weil in der Regel billigere Variante. Häufige Strategien zur Verminderung der Eintrittswahrscheinlichkeit von Risiken in der Software-Entwicklung sind Qualitätssicherungsmaßnahmen und der Bau von Prototypen.

Beispiel:

Das Risiko in einem Software-Projekt, dass der zentrale Repository-Server ausfällt, und somit die Arbeit der Entwickler unterbrochen wird, kann durch eine redundante Infrastruktur mit mehreren synchronisierten Servern und RAID-Systemen gemindert werden.

Kosten: Die Kosten für diese Maßnahme sind offensichtlich die Kosten für die zusätzliche Infrastruktur sowie deren Administration.

Risiko tragen

Es gibt Risiken, gegen die man keine sinnvollen oder bezahlbaren Strategien entwickeln kann. Ebenso gibt es Risiken, deren Risikowert so gering ist, dass sie toleriert werden können. Diese

Beispiel:

Das Risiko in einem Software-Projekt, dass der zentrale Repository-Server ausfällt, kann auch einfach akzeptiert werden (der Server ist ja noch ganz neu, das Repository oder auch das Projekt sind nicht so kritisch).

Kosten: In diesem Fall fallen keine Kosten an, da auch keine Maßnahme umgesetzt wird.

Risiken sollten als Restrisiken gekennzeichnet werden, man akzeptiert, dass diese Risiken eintreten können: Trotzdem können zu diesen Risiken noch Pläne entwickelt werden, wie man auf das Eintreten des Risikos reagiert.

Welche dieser Strategien man anwendet, ist von Risiko zu Risiko unterschiedlich und abhängig von den Erfolgsaussichten und den Kosten der Maßnahme. Bei besonders gravierenden Risiken wird man mehrere Strategien anwenden.

Beispiel:

Das Risiko in einem Software-Projekt, dass der zentrale Repository-Server ausfällt, bedeutet beispielsweise, dass 30 Entwickler mit einem Tagessatz pro Person von 750 Euro zwei Tage lang nicht arbeiten können. Somit beträgt die Schadenshöhe $30 \times 750 \text{ EUR/Tag} \times 2 \text{ Tage} = 45\,000 \text{ EUR}$.

Die Eintrittswahrscheinlichkeit sei mit 10 % angenommen.

Der Risikowert beträgt also $10\% \times 45\,000 \text{ EUR} = 4\,500 \text{ EUR}$.

Setzt man einen zweiten Server ein, vermindert man die Eintrittswahrscheinlichkeit. Die Wahrscheinlichkeit, dass beide Server gleichzeitig ausfallen, beträgt dann $10\% \times 10\% = 1\%$, somit verringert sich der Risikowert um 4 050 EUR auf 450 EUR.

Schlägt nun die Bereitstellung eines zweiten Servers inkl. Administrationsarbeiten mit 3 000 EUR zu Buche, so ist es wirtschaftlich, das Risiko durch die Bereitstellung eines zweiten Servers zu mindern.

Die Gegenmaßnahmen können auch wieder im Brainstorming-Verfahren gefunden werden. Zu jedem Risiko können so eine Menge an Gegenmaßnahmen gefunden werden, die das Risiko vermeiden, transferieren oder mindern.

Zu jedem Risiko hat man jetzt eine ganze Reihe von Gegenmaßnahmen, die umgesetzt werden können. Wichtig ist, alle Maßnahmen zu bewerten. Jede umgesetzte Maßnahme hat einen Einfluss auf die Projektplanung, die Durchführung der Gegenmaßnahmen kann zu höheren Kosten und Terminverschiebungen führen. Außerdem können sich Auswirkungen durch die Umsetzung dieser Maßnahme auf andere Risiken oder Gegenmaßnahmen ergeben. Diese Beziehungen müssen geklärt sein, bevor über die Umsetzung von Maßnahmen entschieden wird.

In jedem Fall sollte bei jeder Maßnahme abgeschätzt werden, wie das Restrisiko nach einer erfolgreichen Umsetzung der Maßnahme einzuschätzen ist.

Diese Verringerung des Risikowerts ist dann den Kosten für die Maßnahme gegenüberzustellen. Eine Maßnahme, die mehr kostet, als sie den Risikowert senkt, sollte kritisch hinterfragt werden. In dem Fall könnte es sogar günstiger sein, das Risiko zu tragen, und keine Maßnahme in Angriff zu nehmen.

Gibt es mehrere mögliche Alternativen, so kann durch die Maßnahmenbewertung die günstigste Maßnahme gefunden werden, die das Risiko am effizientesten senkt.

Abschließend muss noch geklärt werden, ob das verbleibende Restrisiko akzeptabel ist oder ob weitere Maßnahmen ergriffen werden müssen.

6 Maßnahmen umsetzen

Wenn man zu jedem Risiko, das behandelt und nicht einfach nur akzeptiert werden soll, Maßnahmen geplant hat und die Restrisiken definiert sind, können die geplanten Maßnahmen umgesetzt werden. Wichtig ist dabei, dass für jede Maßnahme ein Verantwortlicher definiert ist, der die Durchführung dieser Maßnahme überwacht.

Die Durchführung dieser Maßnahmen muss auch den Projektplan ändern, denn in aller Regel ist sie mit Aufwand verbunden.

7 Risiko- und Maßnahmencontrolling

Die Risiken und die Umsetzung der Gegenmaßnahmen müssen regelmäßig überwacht werden. Die eigentliche Durchführung der Maßnahmen zur Risikominderung kann dabei über den Projektplan überwacht werden. Im Rahmen des Risikomanagements ist dabei insbesondere das Augenmerk darauf zu richten, ob die Maßnahmen tatsächlich auch das Risiko wie beabsichtigt wirksam mindern. Sonst ist der Bewertungs- und Planungsprozess erneut durchzuführen.

Zudem müssen alle Risiken regelmäßig erneut betrachtet werden, da es im Laufe des Projektes, z. B. durch durchgeführte Gegenmaßnahmen zu Änderungen an Eintrittswahrscheinlichkeit und Schadenshöhe kommen kann. Auch können in Abhängigkeit von der Phase im Lebenszyklus oder durch äußere Umstände neue Risiken hinzukommen oder Risiken ganz verschwinden.

Beispiel:

Das Risiko „Die Testtools stehen nicht rechtzeitig zur Verfügung, und dadurch kommt es zu Verzögerungen und ‚Budget overrun‘“ kann in dem Moment aus der Beobachtung entlassen werden, in dem die Testtools tatsächlich installiert sind.

Bei den Maßnahmen ist zu kontrollieren, ob sie durchgeführt wurden, außerdem sollte bewertet werden, ob die beabsichtigten Auswirkungen auf das Risiko (z. B. Senkung der Eintrittswahrscheinlichkeit) eingetreten sind. Sind die Maßnahmen nicht wirksam, so müssen andere Maßnahmen ergriffen werden.

Zusammenfassung

In den vorangegangenen Abschnitten wurden die verschiedenen Aspekte des Risikomanagementprozesses näher beleuchtet.

Nun stellt sich die Frage, wie häufig der Risikomanagement-Zyklus durchgeführt werden sollte? Da auch Risikomanagement Aufwand erfordert, sollte man, wie bei allen Projektmanagement-Disziplinen angemessen vorgehen, und so wenig wie möglich machen. Aber auch nicht weniger als nötig.

Wie viel Risikomanagement?

- * mindestens einmal vor Projektbeginn
- * mindestens einmal während der Projektinitialisierung
- * im Rahmen jedes weiteren Planungszyklus
- * immer dann, wenn Sie steuernd in das Projektgeschehen eingreifen müssen

[Gernert]

Ihre Aufgabe als Projektleiter ist es, Entscheidungen zu treffen, die das Projekt zum Erfolg werden lassen. Ohne ein systematisches Risikomanagement besitzen Sie aber gar keine ausreichende Grundlage, um die richtigen Entscheidungen treffen zu können. Vor Projektbeginn bzw. während der Initialisierungsphase können Sie ggf. noch einige Risiken durch Ihre Planung völlig vermeiden, und andere möglicherweise deutlich vermindern. Solange Sie die Risiken kennen, und es noch Risiken sind, d. h. sie sind noch nicht eingetreten und dadurch zum Problem geworden, können Sie noch aktiv agieren. Im anderen Fall bleibt Ihnen nur noch das Reagieren und ggf. das Krisenmanagement.

„using risk to determine how much planning is enough“

[BoehmTurner]

Risiken bedrohen laufend Ihr Projekt. Zu den wichtigsten Aufgaben eines jeden Projektmanagers gehört es daher, die Risiken zu kennen und in der Projektplanung und -steuerung zu berücksichtigen.

[Gernert] Christiane Gernert, „Agiles Projektmanagement“, Hanser Verlag 2003

[HarrantHemrich] Horst Harrant, Angela Hemmrich „Risikomanagement in Projekten“, Hanser Verlag 2004

[PatzakRattay] Gerold Patzak, Günther Rattay, „Projektmanagement“, Linde-Verlag 2004

[PMF] Deutsche Gesellschaft für Projektmanagement e. V. „Projektmanagementfachmann/-fachfrau“, RKW-Verlag 2004

[Thaller] Georg Erwin Thaller, „Dra-chentöter“, Heise Verlag 2004



Ulrich Andree
Tel.: 0211 9449-5046
E-Mail: ulrich.andree@lds.nrw.de



Dr. Jan Mütter
Tel.: 0211 9449-5434
E-Mail: jan.muetter@lds.nrw.de

Literatur

[BoehmTurner] Barry Boehm, Richard Turner, „Balancing Agility and Discipline“, Addison-Wesley 2003

[DeMarcoLister] Tom DeMarco, Timothy Lister, „Bärentango“, Hanser Verlag 2003



Landesdatenbank NRW Online.
Landesdatenbank **NRW**. Der Internetzugang zu Daten für alle Gemeinden und Kreise Nordrhein-Westfalens

Die Landesdatenbank NRW Online bietet einen umfangreichen und aktuellen Querschnitt aus den wichtigsten Bereichen der amtlichen Statistik und damit die Möglichkeit, wirtschaftliche und soziale Fakten via Internet zu recherchieren und als Tabellen abzurufen.

Enthalten sind Daten über:

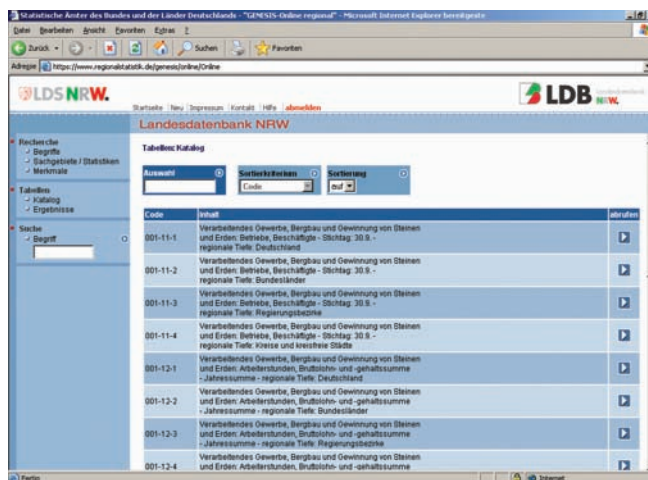
- Gebiet und Bevölkerung
- Gesundheitswesen
- Bildung
- Wahlen
- Erwerbstätigkeit
- Unternehmen und Arbeitsstätten
- Produzierendes Gewerbe
- Bautätigkeit und Wohnungswesen
- Handel und Gastgewerbe
- Verkehr
- Insolvenzen
- Sozialleistungen
- Öffentliche Finanzen
- Preise
- Volkswirtschaftliche Gesamtrechnungen
- Umwelt

Für marktorientierte Unternehmensbereiche, Verwaltungen, Wissenschaft und Forschung erschließen sich wichtige Grundlagen zur Analyse und Entscheidungsfindung.

Bürgerinnen und Bürger erhalten die Möglichkeit, sich umfassend und genau über Fakten zu informieren, die den aktuellen Diskussionen zugrunde liegen.

Zugang zur Landesdatenbank NRW Online

Recherchen in der Landesdatenbank Online sind über eine Stichwort-Suche oder hierarchisch über Sachgebiete möglich. Dazu gibt es variabel gestaltbare Tabellen, d. h. für bestimmte Tabellenpositionen können Merkmale ausgewählt und Abrufe gestartet werden. Eine schnelle Vorschau-Funktion verschafft zuvor einen Eindruck davon, welches Aussehen und welchen Umfang der Abruf einer Tabelle hat. Die Ergebnisse werden nicht nur als HTML-Tabellen angezeigt, sondern es ist auch ein Download im Excel-, CSV- oder HTML-Format möglich. Statistiken, Merkmale und deren Aus-



prägungen werden ausführlich methodisch beschrieben bzw. erläutert, wodurch eine korrekte Interpretation erleichtert wird.

Kontakt

Jörg Mühlenhaupt
Telefon: 0211 9449-4409
joerg.muehlenhaupt@lds.nrw.de
Weitere Informationen finden Sie unter:
<http://www.landesdatenbank-nrw.de/>

Projektmanagement in der Praxis oder: Wo steht mein Projekt?

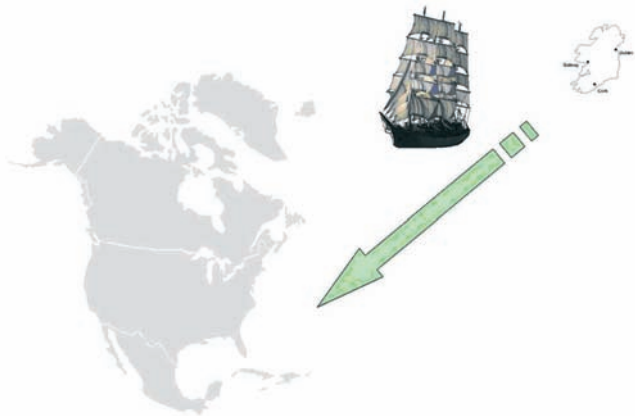


Abb. 1: Der Projektplan

Startet man mit einem Schiff von Irland nach Amerika, so ist es unwahrscheinlich, den anvisierten Hafen mit dem einmal eingestellten Kurs und ohne Gegensteuern zu erreichen.

Genauso ergeht es auch dem Projektleiter eines Softwareprojektes.

Im Projektplan steht der einzuschlagende Kurs, die Realität ist aber erfahrungsgemäß anders als der Plan, manchmal mehr, manchmal weniger.

Im Projektgeschäft, wie auch auf hoher See, ist es daher unerlässlich jederzeit neben dem geplanten Kurs die aktuelle Position zu kennen.

Dabei ist die Positionsbestimmung selbst immer nur Mittel zum Zweck. Es geht einzig und allein darum, das Ziel zu erreichen.

„Prognosen sind schwierig, besonders wenn sie die Zukunft betreffen“, ein Zitat, das u. a. Karl Valentin, Mark Twain und Winston Churchill zugeschrieben wird.

Wenn dies der Fall ist, sollte der Projektleiter genau die Parameter messen, deren Kenntnis er benötigt, um Kursabweichungen zu erkennen. Da die Durchführung von Messungen im Projekt auch Aufwand verursacht, sollte so wenig Zeit wie möglich dafür ver(sch)wendet werden.

Im Verlauf des Beitrags werden wir versuchen, genau die Frage, welche Parameter sinnvollerweise im Projekt gemessen werden sollten, näher zu betrachten und die Aussagen mit eigenen Projekterfahrungen zu unterlegen.

Warum Planung und Fortschrittskontrolle?

Haben Sie Probleme mit der Planung oder Fortschrittskontrolle in Ihrem Softwareprojekt? Falls ja, seien Sie beruhigt, dies ist völlig normal.

Haben Sie schon einmal die Idee gehabt, das alles einfach wegzulassen?

Hier ein Paar Beispiele, warum Planung und Fortschrittskontrolle durchaus notwendig sind:

Am Anfang macht das Projekt gute Fortschritte, Sie sind als Projektleiter zufrieden mit Ihrem Team.

Um bei unserem Bild vom Anfang zu bleiben, das Schiff fährt auf dem richtigen Kurs, der Steuermann hat das Ruder im Griff und den Kompass im Blick, man ist zur geplanten Zeit am geplanten Ort und der Wind bläst aus der richtigen Richtung.

Ähnlich ist es im Softwareprojekt. Hier erfolgt die Fertigstellung der Aufgaben im Statusmeeting immer schön regelmäßig, genau wie geplant. Sie sind als Projektleiter zufrieden und sehen auch keinen Grund, dass irgendetwas Sie und Ihre Mannschaft vom Erfolg des Projektes abhalten könnte.

Irgendwann im Projektverlauf kommen dann aber die Erfolgsmeldungen nicht mehr wie erwartet, sondern etwa Folgendes:

„Die Entwicklung hat da noch ein kleines Problem zu lösen, das ist aber nächste Woche fertig.“ Beim nächsten Statusmeeting erfährt man dann, dass leider doch noch an einer Stelle ein weiteres Modul betroffen ist. Nach einigen Wochen stellt sich heraus, dass die Architektur der Anwendung geändert werden muss, weil doch ein bestimmtes Datenfeld so nicht im Zugriff ist und keiner je daran gedacht hat.

Kennen Sie das?

Jetzt gehen die Probleme erst richtig los: Der Projektplan fällt in sich zusammen. Keiner hat mehr Zeit für die Dinge, die ursprünglich geplant waren, und Sie können als Projektleiter nicht mehr so richtig sagen, wie fertig die ganze Software wirklich ist. Wenn der Auftraggeber das erfährt, könnte sogar das ganze Projekt vorzeitig beendet werden.

Um in dieser Situation eine Aussage treffen zu können, wie viel im Projekt bereits fertig ist und was noch zu tun ist, und um sich selbst, das Projektteam und letztendlich den Auftraggeber zu beruhigen, benötigt der Projektleiter einen Plan und eine Fortschrittskontrolle.

Um einen Fortschritt kontrollieren zu können, wird irgendetwas benötigt, woran man messen kann, die Planung.

Im Beispiel mit der Atlantiküberquerung ist eine Seekarte mit dem eingezeichneten Kurs als Plan nötig, sonst kann der Kurs unterwegs nicht kontrolliert werden. Als Fortschrittskontrolle dienen beispielsweise die regelmäßige Messung von Richtung und Geschwindigkeit durch den Steuermann und der Eintrag des Ist-Wertes in die Seekarte.

Im Softwareprojekt erstellt man einen Plan mit all den Aufgaben, die im Projekt erledigt werden müssen, und den geschätzten Zeiten, die die Durchführung dieser Tätigkeiten vermutlich erfordert, und misst den erreichten Fortschritt.

Planung

In diesem Beitrag soll nicht auf die eigentliche Tätigkeit der Planung von Softwareprojekten eingegangen werden, darüber findet man in der Literatur ausreichend Informationen. Wohl aber werden die Probleme, die in der Praxis mit der Steuerung und Fortschrittskontrolle auftreten, analysiert.

Ist ein Projektplan in ausreichender Detaillierung gegeben, so sind damit

die Soll-Werte vorgegeben, denen man jetzt den Fortschritt, der erreicht wurde, gegenüberstellen muss.

Fortschrittskontrolle

Wie kann aber eine Fortschrittskontrolle in einem Softwareprojekt sinnvollerweise aussehen, um die Aussage, wo das Projekt steht, zu jedem Zeitpunkt des Projekts machen zu können?

Ein Meilenstein versus kleine Aufgaben

Es gibt eine große Bandbreite an Ansätzen, den Ist-Zustand eines Projekts zu charakterisieren. Eine Möglichkeit wäre dabei, die Fertigstellung des Projekts an dem Erreichen eines großen Meilensteins (z. B. „die Software ist einsatzfähig“) zu messen, d. h. es wird nur noch gemessen, ob der Meilenstein erreicht ist oder nicht.

Damit lässt sich über die virtuelle Schifftour nur noch die Aussage treffen, ob wir nach der geplanten Fahrzeit schon in Amerika angekommen sind oder aber nicht. Das ist weder für die Schifffahrt eine angemessene Möglichkeit noch für ein Softwareprojekt.

Wie begegnet man diesem Problem?

Nun, beim Schifffahrt-Beispiel sollte z. B. das Vorbeifahren an einer Insel Anlass sein, die Navigationskarte zu prüfen und den Abgleich von Kompass und Geschwindigkeitsmesser durchzuführen. Ist das Schiff nicht nach der geplanten Zeit an der Insel, stimmt etwas mit den Navigationsmethoden nicht.

Wir zerlegen also die gesamte Schiffsroute in Teilstrecken, deren Erreichen wir überprüfen können.

Was haben wir getan?

Wir haben vor dem Erreichen des großen Meilensteins am Ende des Projekts (im Beispiel also das Erreichen von

Amerika) mehrere kleine Meilensteine definiert (Vorbeifahrt an bestimmten Inseln), die sich während der Projektlaufzeit als Prüfgrößen eignen.



Abb. 2: Der verfeinerte Projektplan

Man ist also gut beraten, mehrere Meilensteine auf dem Weg zum Ziel festzulegen, die gut zu überprüfen sind, um eine Fortschrittskontrolle in einem Softwareprojekt zu ermöglichen.

Inhalte der Meilensteine

Hierbei ist aber gerade die Beschreibung der Meilensteine nicht so einfach.

Zurück zur Schiffsreise: Woher weiß man, dass es die richtige Insel ist? Wann genau fährt man eigentlich an einer Insel „vorüber“? Muss rechts oder links vorbei gefahren werden? Wie groß darf der Abstand sein?

Kurz: Welche Genauigkeit reicht aus, um die Messinstrumente zu eichen?

Ist das nicht festgelegt, können die Abweichungen im Fahrweg recht schnell beachtlich werden. Diese addieren sich dann möglicherweise während der ganzen Fahrt so sehr, dass sich das Ziel gar nicht mehr erreichen lässt. Die exakte Festlegung eines Meilensteins hat somit eine recht hohe Bedeutung.

Wie oft haben Sie schon ein Programm fertig gemeldet bekommen, das dann doch nicht fertig war?

Woran lag das?

Aussagen dazu gibt es viele: z. B. die Programmvorgabe war nicht korrekt und musste nachgebessert werden oder

es war einfach noch ein Fehler im Programm.

Wie auch immer, der angeblich erreichte Meilenstein ist nun plötzlich wieder offen, und der Plan ist schon wieder durcheinander geraten

Exakte Beschreibung der Meilensteine

Begegnen kann man dieser Tatsache nur durch eine möglichst exakte Beschreibung des Meilensteins und der Tätigkeiten, die zu ihm führen.

Hilfreich ist die Einführung von geeigneten Qualitätssicherungsmaßnahmen. Im angenommenen Beispiel könnte man vor der Atlantiküberquerung das genaue Ausmessen der Geschwindigkeit in bekannten Gewässern durchführen, um den Geschwindigkeitsmesser zu eichen.

In der Softwareentwicklung hat sich die Definition von Ein- und Ausgangskriterien für die entsprechenden Tätigkeiten bewährt, die dann jeweils auch geprüft werden müssen, bevor der Meilenstein als erreicht erklärt wird.

Zusammengefasst ist es also wichtig, möglichst alle im Projekt anfallenden Tätigkeiten für alle Beteiligten genau genug beschrieben zu haben, damit das Erreichen dieser Aufgabe auch definitiv erfüllt worden ist, bevor der Meilenstein als erreicht gekennzeichnet wird.

Wie bereits oben erwähnt, wird in diesem Beitrag nicht auf die Probleme einer Planung eingegangen. Trotzdem sei hier noch einmal auf die Angemessenheit der Planung hingewiesen. Da Änderungen der Normalzustand im täglichen Projektgeschäft sind, macht es keinen Sinn, zum Projektbeginn einen vollständigen Feinplan aufzubauen, sondern man wird sich zunächst mit einer Grobplanung begnügen, die eine für die anstehenden Entscheidungen

adäquate Informationsdichte enthält (siehe [vonHagenMütter1] und [BoehmTurner]).

„Planen auf Vorrat ist wenig sinnvoll. Beginnen Sie daher immer erst dann mit der Feinplanung, wenn die jeweilige Aufgabe konkret ansteht.“
[Gernert]

Wie weit müssen die Arbeitsschritte verfeinert werden?

Die Kunst in der Softwareentwicklung liegt nun darin, eine angemessene Verfeinerung der Arbeitsschritte zu finden.

Wenn die Arbeitsschritte zu lang sind, wird der Zustand nur einmal während der Laufzeit gemessen, nämlich nach der abgelaufenen Plandauer. In diesem Fall wird möglicherweise zu spät festgestellt, dass das Projekt diesen Meilenstein nicht erreicht hat.

Auf der anderen Seite verursacht die Positionsbestimmung auch Aufwand, und im Extremfall werden die Mitarbeiterinnen und Mitarbeiter im Projekt vollständig mit dem Erstellen von Fortschrittsberichten ausgelastet sein. Es ist einleuchtend, dass eine 100 %ige Auslastung der Mitarbeiterinnen und Mitarbeiter durch die zur Positionsbestimmung nötigen Aufgaben das eigentliche Projektziel, das Projekt erfolgreich abzuschließen, keinen Schritt näher bringt.

Wie aber findet man eine angemessene Planungs- und Berichtsebene? Sollen die Beschäftigten auf fünf Minuten genau ihre Tätigkeiten erfassen, oder genügt es das tageweise zu tun? Und wie oft müssen Sie als Projektleiter diese Informationen haben?

Im weiteren Verlauf wird dieser Punkt noch einmal aufgegriffen.

Die Interpretation der Fortschrittsdaten. Sind wir noch auf dem richtigen Weg?

Ein weiteres Problem könnte sich auf der Beispielstour möglicherweise so gestalten:

Der Steuermann arbeitet immer mit dem gleichen Kompass und Geschwindigkeitsmesser. Zeigt eines dieser Instrumente falsch an, begeht er unwissentlich einen Fehler. Sie gehen auch als Kapitän, und nach ihren Informationen durchaus logisch und richtig, davon aus, dass Sie keinen Fehler machen. Sie glauben zu wissen, dass Sie genau an der Stelle des Meeres sind, die Ihre Navigations-Karte angibt. Merken werden Sie Ihren Irrtum erst, wenn nach ihren Daten die Zielküste in Sicht kommen müsste, dies aber nicht so ist.



Abb. 3: Der unerwartete Fehler

In der Projektleitung von größeren Softwareprojekten sollte man sich bewusst sein, dass nicht jede Information, auch wenn sie gut gemeint ist, richtig ist. Es ist die Aufgabe des Projektleiters geeignete Maßnahmen zu ergreifen, um das zu erkennen.

So, wie sich der Steuermann niemals ausschließlich nach dem Kompass richten wird, sondern immer auch noch versuchen wird, den Kurs anhand von Sonne und Sternen oder sogar mit Hilfe von Landmarken zu bestimmen, sollte der Projektleiter mehrere Metriken einsetzen, um die verschiedenen Informationen plausibilisieren zu können. Beispielsweise ist die Aussage „Die Software ist zu über 90 % fertig“ zu hinterfragen, wenn die Zahl der offenen Fehler und Änderungsanträge noch recht hoch ist.

Die ungeplanten Tätigkeiten

Gerade in der Softwareentwicklung ist es völlig normal, dass irgendetwas nicht bedacht wurde. Projekte in der Softwareentwicklung zeichnen sich dadurch aus, dass hier eine völlig neue Software entsteht. Daher ist es auch nur logisch, dass kein Fachkonzept fehlerlos erstellt wird, kein Architekt an alles vorher denken kann und kein Entwickler fehlerfrei arbeitet.

Eine geeignete Maßnahme, diesen Umstand zu berücksichtigen, ist die Einplanung von Pufferzeiten. Das ist auch in unserem Beispiel der Schiffstour möglich. Man rechnet die längste mögliche Fahrzeit aus (Funktion von Geschwindigkeit, Wind, Strömung etc.), addiert noch etwa 5 % dazu und schreibt den Wert in den Fahrplan. Die Kunden werden zufrieden sein, denn das Schiff kommt mit seiner Ladung dann öfter pünktlich an, als wenn man nur die reine Fahrzeit angegeben hätte.

Es macht also durchaus Sinn, auch in einem gut durchgeplanten Projekt einen Puffer einzuplanen. Dies kann dadurch geschehen, dass man die Anzahl der benötigten Personalressourcen entsprechend erhöht oder die Projektlaufzeit verlängert.

Es gibt wenige größere Projekte, in denen nichts Unvorhergesehenes den Projektplan verändert. In jedem Projekt, in dem Tools benutzt werden, wird es auch immer jemand geben müssen, der sich z. B. um das Upgrade dieser Tools kümmert. Dieses im Detail einzuplanen, ist gerade bei langlaufenden Projekten so gut wie unmöglich. Daher empfiehlt sich dieser Puffer. Die Größe des Puffers ist davon abhängig, wie hoch die Risiken im Projekt sind. Ist z. B. die Entwicklerteams neu zusammengestellt worden, muss der Puffer deutlich höher angesetzt werden, als wenn mit den Beteiligten schon mehrere Programme erfolgreich erstellt wurden. Gleiches gilt für die benutzte Technologie, ist

sie neu, ist der Pufferwert hoch anzusetzen, ist sie bekannt, kann der Wert kleiner sein. Dieser Bereich der Softwareplanung ist sehr komplex und soll hier nicht weiter vertiefend beleuchtet werden. Wichtig ist nur, die Pufferzeiten nicht zu vernachlässigen.

Kosten von Planung und Fortschrittskontrolle

Planung und Fortschrittskontrolle kosten Geld. Andererseits kann eine vernachlässigte Planung und Fortschrittskontrolle möglicherweise noch viel mehr Geld kosten.

Der im Beispiel bereits bemühte Kapitän wird auf seinem Schiff einen Steuermann mitnehmen, der sich auf der ganzen Reise nur um das Gegensteuern und das Kurshalten zu kümmern hat. Er wird gerade auf älteren Schiffen einen Navigator anheuern wollen, der ausschließlich für die Positionsbestimmung und Kursberechnung zuständig ist.

Es ist unbestritten, dass die zusätzlich benötigten Seeleute Geld kosten, aber was würde es kosten, wenn die Ladung verspätet oder gar nicht an ihr Ziel gebracht würde?

Ähnlich verhält es sich in einem Softwareprojekt.

Dazu einige Daten aus einem derzeitigen LDS-Projekt:

In diesem Projekt müssen alle im Team ihre Zeiten wöchentlich an die Teilprojektleiter berichten, die wiederum die aggregierten Zahlen an den Gesamtprojektleiter weitergeben.

Jede(r) der etwa 60 Mitarbeiter/-innen benötigt dafür wöchentlich jeweils etwa 10 Minuten. Die 5 Teilprojektleiter benötigen zudem jeweils 2 Stunden wöchentlich um ihre Teilprojektpläne zu aktualisieren, neu zu planen und die aggregierten Fortschrittszahlen an

den Gesamtprojektleiter zu berichten. Die Fortschrittsbestimmung im Projekt nimmt also maximal 20 Personenstunden wöchentlich in Anspruch.

Andererseits wird in einem Projekt dieser Größe auch wöchentlich ein Aufwand von 60 Personenwochen verbraucht, für den Kosten entstehen. Läuft das Projekt nur eine Woche in die falsche Richtung, können leicht Kosten für mehr als ein Personenjahr verschwendet werden.

Wenn man sich den Gesamtverbrauch bewusst macht, ist der Wert für die Kontrolle gering. Es ist also durchaus zu rechtfertigen, diesen Aufwand für die Datensammlung als Risikominimierung aufzuwenden.

Risikomanagement

Stellt man erst nach einem Monat oder noch später fest, dass in die falsche Richtung entwickelt wurde, sind dem Projekt bereits erhebliche Kosten entstanden, ohne den damit geplanten Fortschritt zu erzielen. Das Risiko, dass dieses Projekt die geplanten Ziele nicht mehr erreichen kann, ist dann sehr groß.

Hat der Kapitän mit seinem Schiff auf dem Weg von Irland nach Amerika (ohne Navigator und Steuermann) erst einmal den Äquator passiert, ist die Wahrscheinlichkeit sehr groß, dass er und seine Mannschaft das angepeilte Ziel mit den mitgenommenen Trinkwasservorräten nicht mehr erreichen können.

Je knapper seine Trinkwasservorräte anfangs bemessen sind, desto weniger darf er vom geplanten Kurs abweichen. Das Risiko, dass zum Ende hin die Vorräte nicht ausreichen, ist dann sehr groß. Offenkundig kann der Kapitän dieses Risiko in gewissen Grenzen sehr einfach mindern, indem er mehr Wasser mitnimmt, als nach Plan unbedingt erforderlich ist, er also einen Puffer für Kursabweichungen berücksichtigt.

Ist auf der anderen Seite der Vorrat nicht übermäßig reichlich bemessen, das Risiko also groß, so wird er seine Position sehr häufig bestimmen, um Kursabweichungen möglichst frühzeitig erkennen und gegensteuern zu können.

Dieses Beispiel macht deutlich, dass der Aufwand, der für die Positionsbestimmung auch im LDS-Projekt als angemessen angesehen werden kann, von den das Projekt bedrohenden Risiken abhängt.¹⁾

In termin- und budgetkritischen Projekten ist eine kürzere Frequenz für die Fortschrittskontrolle und damit ein höherer Aufwand einzuplanen als in weniger kritischen Projekten.

Die Maßnahmen zur Reduzierung der Risiken müssen in einem angemessenen Verhältnis zu den Risiken selbst stehen.

Methoden

In der Literatur findet man viele Aussagen darüber, wie das Projektcontrolling, d. h. die Projektüberwachung und der Vergleich der Ist-Ergebnisse mit dem Projektplan, durchzuführen ist.

Die zentralen Parameter zur Beschreibung des Projektzustands sind Leistungsumfang/Qualität des zu erstellen Produkts, Termin und Budget.

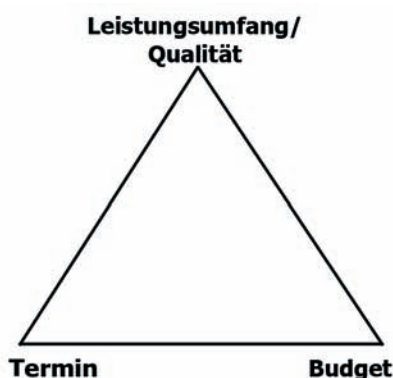


Abb. 4: Magisches Dreieck des Projektmanagements

1) Siehe hierzu im vorliegenden Heft (Seite 28 ff.) den Beitrag von Ulrich Andree und Jan Mütter „Risikomanagement in Softwareprojekten“.

Häufig werden diese in enger Wechselwirkung untereinander stehenden Parameter im „magischen Dreieck des Projektmanagements“ (siehe [PMF], Kap.1.6) dargestellt.

In dem hier vorliegenden Beitrag soll der Parameter Leistungsumfang/Qualität des zu erstellenden Produkts bzw. der Dienstleistung nicht weiter betrachtet werden. Es wird dabei angenommen, dass der Leistungsumfang hinreichend eindeutig festgelegt ist und die Umsetzung regelmäßig von den Prozessen Qualitätssicherung und Änderungsmanagement verfolgt wird.

Wo stehen wir?

Im Projektplan wurde festgelegt, welche Aktivitäten wann und von wem durchzuführen sind und wie viel Ressourcen (Aufwand) dafür zur Verfügung stehen (siehe auch [vonHagen Mütter2]).

Läuft nun das Projekt, so muss der Projektleiter ständig den Überblick darüber behalten, ob alles noch „nach Plan“ abläuft oder ob er „gegensteuern“ muss, um das Projekt wieder auf Kurs zu bringen.

Aufwandsmessungen

Ein zentraler Projektparameter ist der Aufwand der einzelnen Aktivitäten. Dieser Parameter bestimmt in vielen Projekten maßgeblich das Budget.

Um den Überblick über den Projektstatus zu behalten, ist es daher erforderlich, den Aufwand für die Durchführung der geplanten Aktivitäten zu beobachten. Dazu werden in regelmäßigen und nicht zu großen Abständen die tatsächlich entstandenen Aufwände gemessen.

Dabei sollten möglichst die Aufwände für die Aktivitäten gemessen werden, die auch im Projektplan beschrieben sind.

Ein Arbeitspaket (AP) ist ein Teil des Projekts, der im Projektstrukturplan nicht weiter aufgegliedert ist und auf einer beliebigen Gliederungsebene liegen kann (DIN 69901, [PMF])

Misst man die Aufwände auf einer größeren Skala, fasst also geplante Aktivitäten zu größeren Einheiten zusammen, so ist es anschließend nicht mehr möglich, den aktuellen Status einer einzelnen Aktivität zu bestimmen. Im umgekehrten Fall einer Aufwandsmessung auf einer feineren Skala als jener der im Projektplan dargestellten Aktivitäten spiegeln die Messwerte eine Genauigkeit vor, die sich im Projektplan nicht wiederfindet. In diesem Fall wäre der Aufwand für die Aufwandsmessung und das Projektcontrolling zu hoch.

Im nächsten Schritt kann dann der entstandene Gesamtaufwand (Ist-Aufwand) mit dem Aufwand verglichen werden, der nach Plan hätte entstehen sollen (aktueller Plan-Aufwand).

Diese Information lässt nur eine sehr grobe Aussage über den Projektstatus zu, nämlich die, ob der bisherige Gesamtaufwand wie geplant oder höher oder geringer ist. Eine Grundlage für fundierte Entscheidungen ist dies jedoch normalerweise noch nicht.

Eine weitere Verfeinerung der Statusbeurteilung ergibt sich, wenn neben dem tatsächlich entstandenen Aufwand (Ist-Aufwand) auch jeweils der noch offene Aufwand (Soll-Aufwand) gemessen und dadurch in regelmäßigen Abständen die der Projektplanung zugrunde liegende Aufwandsschätzung verifiziert bzw. aktualisiert wird.

Die Summe des aktuellen Ist-Aufwands und des aktuell noch offenen Aufwands ergibt dann eine aktualisierte Schätzung für den zu erwarten-

den Gesamtaufwand, der dann wiederum mit dem nach Projektplan zu erwartenden Gesamtaufwand verglichen werden kann.

Dieser Vergleich gibt dem Projektleiter dann einen Hinweis darauf, ob das Projektbudget eingehalten oder überschritten werden wird.

Die praktischen Aspekte der Aufwandsmessungen (Wer berichtet? Wie wird berichtet? Personenstunden oder Personentage? etc.) werden später in diesem Beitrag besprochen.

Fortschrittsmessung

Die oben beschriebene Aufwandsmessung bezieht sich nur auf den Parameter „Budget“. Genauso wichtig ist aber der Parameter „Termin“. Hierzu muss in ebenso regelmäßigen Abständen der aktuell erreichte Fortschritt des Projekts gemessen werden.

Beispielsweise kann ein Projektzustand erwünscht sein, bei dem der Aufwand schneller anfällt als geplant, wenn im Gegenzug auch das Produkt schneller fertig wird.

Zur Beurteilung der Effizienz eines Projekts sind unbedingt Fortschritt und Aufwand in Beziehung zueinander zu setzen.

Bevor auf diesen Punkt näher eingegangen wird, ist zu hinterfragen, was eigentlich genau „Fortschritt“ ist und wie man diesen messen kann.

Unter Fortschritt bzw. Fortschrittsgrad versteht man hier den Grad, zu dem das Produkt fertig gestellt ist.

Auch diese Beschreibung ist noch nicht ganz zufriedenstellend, gibt es doch u. a. folgende zwei Möglichkeiten, dies zu definieren.

Statusschrittmethode

Eine Möglichkeit, den Fortschritt einer Aktivität und damit des gesamten Projekts zu messen ist die sog. Statusschrittmethode (siehe [PMF]). Dabei wird jeder Aktivität eine diskrete Anzahl an möglichen Status zugeordnet, z. B. geplant = 0 %, in Bearbeitung = 50 %, erledigt = 100 %.

Bewertet man nun den Zustand mit diesen Status und bildet anschließend eine mit dem Aufwand der Aktivität gewichtete Summe über alle Aktivitäten, so liefert dieser Wert eine gute Näherung des Fortschrittsgrads.

Es ist auch denkbar, lediglich die beiden Status „noch nicht erledigt“ = 0 % und „erledigt“ = 100 % zuzulassen, insbesondere, wenn das Projekt sehr viele Aktivitäten enthält, so dass sich die „Ungenauigkeiten“ herausmitteln. (Bei einem Projektplan mit 100 ungefähr gleich aufwändigen Aktivitäten trägt jede Aktivität nur mit 1 % zum Gesamtfortschritt bei. Somit ist der „Fehler“, der durch die Bewertung mit nur zwei Zuständen entsteht, maximal 1 %)

Die Beschränkung auf nur zwei Zustände entspricht auch einer Meilensteinmethode, bei der, wie in unserem Schiffs-Beispiel, einfach nur gezählt wird, welche Meilensteine bereits erreicht bzw. welche Inseln bereits passiert wurden.

Aus dem gerade Gesagten wird deutlich, dass es sich in einem Projekt mit wenigen und dafür größeren Aktivitäten empfiehlt, mehrere verschiedene Status zu unterscheiden.

Fortschrittsmessungen über Aufwandsmessungen

Im vorangegangenen Abschnitt sind wir auf die Aufwandsmessungen eingegangen. Wenn man den Ist-Aufwand jeder Aktivität des Projektplans misst und annimmt, dass das Projekt gemäß

dem Plan abläuft, so gibt das Verhältnis des gemessenen Ist-Aufwands zum geplanten Gesamtaufwand eine erste Näherung für den Fortschrittsgrad, die aber keinerlei Hinweis auf Budgetüber- oder -unterschreitungen erlaubt.

Häufig laufen Projekte aber nicht wie geplant. Aktivitäten, die in der Initialisierungsphase noch sehr komplex erschienen, erweisen sich nun als vergleichsweise einfach, andere Aktivitäten, die jetzt offenkundig erforderlich sind, wurden anfangs überhaupt nicht berücksichtigt.

Somit ist es sehr sinnvoll, den zu erwartenden Gesamtaufwand projektbegleitend regelmäßig neu zu schätzen, wie bereits oben beschrieben. Dann liefert das Verhältnis aus dem gemessenen Ist-Aufwand und der Summe aus Ist-Aufwand und dem geschätzten noch offenen Aufwand ein Maß für den Fortschrittsgrad einer jeden Aktivität und des gesamten Projekts.

In der Literatur findet sich noch eine Reihe weiterer Methoden zur Bestimmung des Fortschrittsgrads (siehe z. B. [PMF]).

Was bedeutet dies für ein Projekt?

Hat man einmal die Ist-Werte des Projekts erhoben, d. h. der tatsächlich erreichte Projektfortschritt und der dafür angefallene Aufwand sowie der aktuell noch zu erwartende Aufwand bis zur Fertigstellung sind bekannt, dann kann man sich der alles entscheidenden Frage zuwenden: „Werden wir das gesteckte Ziel erreichen?“

Earned-Value-Analyse

Für die Beantwortung dieser Frage eignet sich besonders die sog. „Earned-Value-Analyse“ (siehe [PMBOK] und [PMF]), die in der Literatur auch als „Methode des Fertigstellungswerts“

oder als „Arbeitswertmethode“ bekannt ist.

Dabei werden die gemessenen Ist-Größen (Aufwand und Fortschrittsgrad) mit den geplanten Größen in Beziehung gesetzt und Hochrechnungen auf den verbleibenden Projektverlauf gebildet.

Die Gesellschaft für Projektmanagement [GPM] und das Project Management Institute [PMI] verwenden dabei für denselben Sachverhalt unterschiedliche Begriffe, die in folgender Übersicht gegenübergestellt sind (siehe auch [PMF] und [PMBOK]).

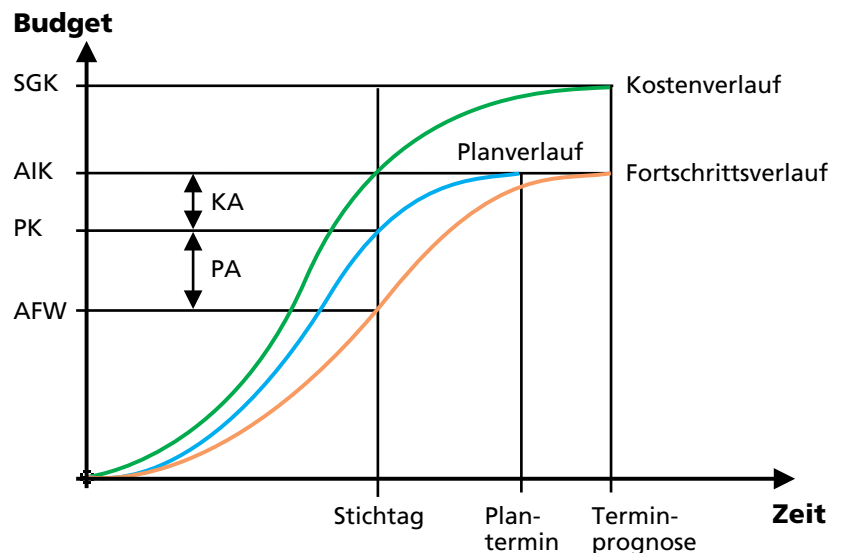


Abb. 5: Typischer Plan-, Kosten- und Fortschrittsverlauf

Gegenüberstellung der Begriffe der Earned-Value-Analyse nach GPM- und PMI-Nomenklatur	
Fertigstellungswert [PMF]	Earned value analysis [PMBOK]
PK Plankosten (Stichtag)	PV planned value total budgeted planned amount to be spent on an activity up to a given point
AIK Aktuelle Ist-Kosten	AC actual cost total cost incurred in accomplishing work on the activity during a given period
AFW Fertigstellungswert Der dem Ist-Fortschrittsgrad (Fertigstellungsgrad (FGR)) entsprechende Anteil der geplanten Gesamtkosten (GPK) $AFW = GPK \cdot FGR \text{ (Ist)}$	EV earned value budgeted amount for the work actually completed
KA Kostenabweichung $KA = PK - AIK$	CV cost variance $CV = EV - AC$
PA Planabweichung $PA = PK - AFW$	SV schedule variance $SV = EV - PV$
ZK Zeitplan-Kennzahl $ZK = PK / AFW$	SPI schedule performance index $CPI = EV / PV$
EF Effizienz-Kennzahl $EF = AFW / AIK$	CPI cost performance index $CPI = EV / AC$
SGK Schätzwert für Gesamtkosten $SGK = GPK / EF$	EAC estimate at completion $EAC = ETC + AC$
	ETC estimate to complete • $ETC = (\text{total PV} - EV \text{ to date})$

Diese Betrachtung der Projektdaten erlaubt projektbegleitend Hochrechnungen für den verbleibenden Projektverlauf durchzuführen, um so schnell und sicher erkennen zu können, ob zum jetzigen Zeitpunkt Korrekturmaßnahmen ergriffen werden müssen, um das Projektziel zu erreichen.

In obiger Abbildung 5 sieht man einen (oft) typischen Verlauf von Kosten und Fortschritt im Projekt: Zum Stichtag, an dem die Messungen durchgeführt werden, sind die inzwischen aufgelaufenen Kosten höher als geplant, während der erreichte Fortschritt hinter den Erwartungen zurückbleibt. Aus der

Earned-Value-Analyse kann man sehr schnell erkennen, dass das Projekt voraussichtlich sowohl den Kosten- als auch den Terminrahmen überschreiten wird.

Die Praxis

Soweit die Theorie, in der Praxis lauern die Fallstricke aber an ganz unvermuteten Stellen.

Was passiert mit den Fortschritts-Daten?

Die Praxis zeigt, dass meistens viele Vorbehalte gegen die Sammlung von Fortschrittsdaten bei den Mitarbeiterinnen und Mitarbeitern vorhanden sind. Hier hilft nur, die anfänglichen Pläne durchzuhalten, darauf zu bestehen, dass die Daten geliefert werden, und bei den Kolleginnen und Kollegen Überzeugungsarbeit zu leisten.

Sind die Widerstände gemeistert und die Datenlieferung sichergestellt, drängen sich weitere Fragen auf, z. B.:

- Wie viele Daten werden benötigt?
- Wie oft benötigt man die Daten?

Es leuchtet auch hier ein, dass zu viele Daten die Verwaltung unmöglich machen und zu wenige Daten keine ver-

nünftige Aussage über den eigentlich erzielten Fortschritt erlauben. Es müssen also für genau so viele Tätigkeiten die verbrauchten Zeiten gesammelt werden, wie es nötig ist, um mit dem Plan vergleichen zu können. Je detaillierter der Plan ist, desto detaillierter muss auch die Aufwandserfassung sein.

Ein Beispiel, das sich in der Praxis bewährt hat:

Man sammelt die verbrauchten Zeiten für genau die im Projektplan aufgenommenen Tätigkeiten. Mindestens die im Gesamtprojektplan aufgenommenen Tätigkeiten gehören auch in die Aufwandsstatusliste als „Projektkostenstelle“. Eine weitere Verfeinerung führt meist zu einem erhöhten Aufwand, ohne einen größeren Nutzen zu haben, kann aber bei großen Teilprojekten nötig werden, damit der Grundsatz der Fortschrittskontrolle in diesen Teilprojekten möglich ist. In diesem Fall sollten dann unbedingt in dem darüber liegenden Gesamtprojekt die Aufgaben im Plan des Gesamtprojekts zusammengefasst werden, und auch nur diese Zusammenfassung sollte berichtet werden.

Bewährt hat sich auch ein Auffangbehälter für sonstige Aufgaben („Projektkostenstelle“ für Sonstiges). Darauf kann bei Bedarf gebucht werden, wenn nicht sicher ist, ob die Aufgabe tatsächlich im Plan angegeben war. Natürlich muss hier kontrolliert werden, dass die darauf anfallenden Stunden nicht ins Unermessliche steigen. Geschieht dies, sollte eine neue geeignete Projektkostenstelle in Absprache mit den Mitarbeiterinnen und Mitarbeitern entstehen. Im Normalfall ist dann der Plan anzupassen, sonst funktioniert der Abgleich nicht.

Beispielsweise wurden in einem LDS-Projekt gute Erfahrungen mit einer Projektkostenstelle für Teammeetings etc. gemacht. Durch solche allgemeine Projektkostenstellen werden die Genauigkeiten bei den tatsächlichen Aufgabekostenstellen höher.

Zwingt man Mitarbeiterinnen bzw. Mitarbeiter dazu, alle Anwesenheitsstunden zu berichten, dann werden oft pro Tag acht Stunden angeschrieben und nicht die tatsächliche Zeit. Wenn die Mitarbeiter/-innen trotzdem so vorgehen, ist der Fehler, der entsteht, nicht so groß. Die Mitarbeiter/-innen wurden ja auch für acht Stunden eingeplant.

Die Beschäftigten sollten wöchentlich die Zahlen sammeln, dann ist die Erinnerung an die Tätigkeiten meist noch so frisch, dass eine Zeiterfassung ohne sehr großen Aufwand möglich ist. Sammeln Sie regelmäßig die Daten und kontrollieren Sie, wer ggf. noch nicht geliefert hat. Dieses Vorgehen wird nach einiger Zeit fast von alleine funktionieren.

Wenn die Daten dann erfasst wurden, kann man damit beispielsweise die oben bereits beschriebene Earned-Value-Methode nutzen, um sich rasch ein genaues Bild vom aktuellen Projektstatus und der Entwicklung des Projektverlaufs zu machen.

Fortschrittsmessung in Projekten mit sich ändernden Anforderungen

Jetzt kommen wir zu einem Problem, dass eigentlich so in einem geplanten Projekt überhaupt nicht vorkommen sollte, in der Praxis allerdings gar nicht so selten ist. In der Softwareentwicklung kommt es immer wieder vor, dass doch eine Tätigkeit z. B. aufgrund von Fehlerbehebungen und Änderungen wiederholt werden muss. Was macht man dann mit seiner Fortschrittsmessung?

Eine Tätigkeit im Projektplan ist bereits als erledigt gekennzeichnet und plötzlich ist ein Korrekturlauf unerlässlich. Dies ist beispielsweise schon der Fall, wenn sich nur ein Fehler in einem Programm eingeschlichen hat, der in einem Test gefunden wurde. Die Programmierung muss nun ein weiteres Mal die Tätigkeit „Programmierung

von Programm xy“ durchführen und einen erneuten Test starten. Schon entsteht ein Darstellungsproblem im Projektplan.

Beide genannten Aufgaben sind streng genommen neue Tätigkeiten, die im Projektplan gar nicht vorhanden sind, obwohl es um die Erstellung des gleichen Programms und den Test wie im Projektplan angegeben geht.

Genau genommen müsste der Projektleiter für diesen Fall den Plan abändern, eine zweite Aufgabe „Überarbeitung des Programms xy“ und „Nachtest des Programms xy“ einarbeiten. Dadurch wird sofort das Projekt-Budget überschritten.

Die Lösung dafür ist entweder die Einplanung eines Puffers, den wir schon beschrieben haben, oder die Möglichkeit, schon bei der Planung mehrere Durchläufe von vornherein vorzusehen.

In diesem Fall können die auf das Programm xy gebuchten Aufwände, z. B. gegen die Planwerte des zweiten Durchgangs oder den Puffer abgerechnet werden. Durch diesen Trick sind Sie weiterhin im Budget und behalten die Kontrolle über das Projekt. Den Überblick werden Sie allerdings sehr schnell verlieren, wenn zu viele Tätigkeiten durch den Puffer abgefangen werden sollen.

Spätestens jetzt sollte neu geplant werden. Verteilen Sie einen Teil des hoffentlich ausreichend vorhandenen Puffers auf die neu angefallenen Tätigkeiten, sonst wissen Sie schon nach kurzer Zeit nicht mehr, wie es um das Projekt steht. Haben Sie keinen Puffer, kann eigentlich nur noch der Weg zum Auftraggeber helfen, der dann das Budget aufstocken muss.

Eine bewährte Methode ist es, mindestens drei Iterationen für die Fehlerbehebung aller Programme mit einzuplanen. Dabei reicht es für die Fortschrittskon-

trolle, wenn die Aufwände für alle diese Fehlerbehebungen in diesem Schritt zusammengefasst werden. Erfahrungsgemäß ist es in diesem Fall – wie beim obigen Korrekturlauf – für die Mitarbeiter/-innen einfacher, auf das betroffene Programm zu buchen. Der Projektleiter sollte dann beim Planabgleich die Werte entsprechend zuordnen.

Eine kleine Anmerkung zu der Anzahl der eingeplanten Iterationen für die Fehlerbehebung. Natürlich ist die Zahl drei nur ausreichend, wenn sich die Anforderungen nicht ändern und die Qualität der vor der Programmierung liegenden Prozesse (Architektur, Design etc.) gut genug ist. Wird im Projekt eine Anforderung geändert, die eine Umprogrammierung zufolge hat, werden auch diese Programme mit mindestens zwei bis drei weiteren Iterationen durch den Test laufen, bis wieder alles fehlerfrei läuft. Der Vorteil für unseren Projektplan und das Budget ist, dass die Änderung vom Auftraggeber über das Änderungsmanagement freigegeben werden muss und somit diese neuen Iterationen in einen neuen Projektplan aufgenommen werden können. Nur mit dieser Vorgehensweise werden Sie auch weiterhin die Aussage treffen können, welchen Fortschritt das Projekt gemacht hat und welche Programmabschnitte offen sind.

Wir sehen also, für eine effektive Fortschrittsplanung gibt es eigentlich keine „wiederauflebenden“ Tätigkeiten. In Wirklichkeit sind das immer neue Tätigkeiten, die ggf. noch nicht im Plan waren. Versäumen Sie es, diese in dem Plan mit aufzunehmen, werden Sie früher oder später den Überblick verlieren. Es hilft dann nur ein Puffer und eine Neuplanung.

Zusammenfassung

Der Kapitän hat auf seiner Atlantikreise alle Ratschläge befolgt. Er verzeichnete auf seiner Seekarte die optimale



Abb. 6: Am Ende des Projekts

Route, gleichwohl berücksichtigte er aber Reserven für Unvorhergesehenes.

Weil die Reise von Irland nach Amerika weit ist, viel Unvorhergesehenes auf dem Ozean lauert und weil dadurch die Mannschaft und die Reise insgesamt bedroht ist, hat er die Puffer großzügig bemessen.

Die Position des Schiffes ließ er mehrmals täglich bestimmen (in Softwareprojekten haben wir gute Erfahrungen mit einer wöchentlichen Berichterstattung gemacht) und falls erforderlich gegensteuern.

So hat er zum Schluss sein Schiff, seine Mannschaft und die Ladung sicher ans Ziel gebracht, so wie das LDS NRW es Ihnen auch für Ihre Softwareprojekte wünscht.

Literatur/Links

[BoehmTurner] Barry Boehm, Richard Turner, „Balancing Agility and Discipline“, Addison-Wesley 2003

[Gernert] Christiane Gernert, „Agiles Projektmanagement“, Hanser Verlag 2003

[GPM] Deutsche Gesellschaft für Projektmanagement www.gpm-ipma.de

[ICB] „IPMA competence baseline“, Wissensspeicher der IPMA: www.gpm-ipma.de; Caupin, G.; Knöpfel, H.; Morris, P.; Motzel, E.; Pannenbäcker, O. (Hrsg.): ICB - IPMA Competence Baseline. 2. überarb. Aufl., IPMA-Verlag, Zürich, 1999, ISBN 3-00-004057-9

[IPMA] International Project Management Association www.ipma.ch

[PatzakRattay] Gerold Patzak, Günther Rattay, „Projektmanagement“, Linde-Verlag 2004

[PMBOK] „Guide to the Project Management Body of Knowledge“, PMI. Wissensspeicher des PMI zum Projektmanagement: www.pmi.org; die Paperback Ausgabe: ISBN: 193069945X

[PMF] Deutsche Gesellschaft für Projektmanagement e.V. „Projektmanagementfachmann/-fachfrau“, RKW-Verlag 2004

[PMI] Project Management Institute - www.pmi.org

[PMK] „PM-Kanon“, GPM – Die „deutsche Version“ der ICB: www.gpm-ipma.de; Motzel, E.; Pannenbäcker, O. (Hrsg.): Projektmanagement-Kanon. Der deutsche Zugang zum Project Management Body of Knowledge. TÜV-Verlag, Köln, 1998, ISBN 3-8249-0498-5

[vonHagenMütter1] Ulrich von Hagen, Jan Mütter, „Agil und extrem (2) – Praktische Erfahrungen mit agilen Methoden“, LDVZ-Nachrichten 1/2005

[vonHagenMütter2] Ulrich von Hagen, Jan Mütter, „Aufwandsschätzungen in Softwareprojekten“, LDVZ-Nachrichten 2/2005



Thomas Reiff (IBM)
Tel.: 0171-222 7457
E-Mail: reiff@de.ibm.com



Dr. Jan Mütter
Tel.: 0211 9449-5434
E-Mail: jan.muetter@lds.nrw.de

Server Based Computing

Auf dem Weg in die Zukunft

Konsolidierung. Ein Schlagwort, das in den letzten Jahren überall im IT-Bereich zu hören ist. Doch was bedeutet Konsolidierung für die einzelnen Behörden und Einrichtungen der Landesverwaltung NRW? Wo bewegt sich die IT-Infrastruktur in den nächsten Jahren hin? Wie bin ich als Verantwortlicher in meiner Behörde in der Lage, unter den gegebenen Rahmenbedingungen meine IT-Infrastruktur so effizient wie möglich zu gestalten?

Unsere Antwort darauf lautet „Server Based Computing“! Der vorliegende Artikel soll die technisch interessierten Leser mit den Eigenschaften und Vorteilen des Server Based Computing vertraut machen. Es handelt sich um eine Schlüsseltechnologie für die Konsolidierung von IT-Infrastruktur. Zunächst werden die Grundlagen und der Einsatz von Server Based Computing im LDS NRW dargestellt. Im Anschluss daran wird aufgezeigt, wie diese Technologie in einer IT-Landschaft genutzt werden kann und welche Hilfestellungen das LDS NRW bei der Einführung und für einen optimalen Einsatz von Server Based Computing anbietet. Letztlich wird sich zeigen, wie Sie mithilfe von Server Based Computing den Weg in die Zukunft erfolgreich bewältigen können.

Vergangenheit und Gegenwart

Wir erinnern uns an die Entwicklung der IT-Infrastruktur in den 1970er- und 1980er-Jahren: Es war die Zeit der Mainframes und Großrechner. Viele der älteren Kollegen haben mit dieser Technik gearbeitet. Im LDS NRW gab es Computeranlagen in Wohnzimmergröße, die wassergekühlt die aufwändigen IT-Verfahren der Statistik berechneten. Ein fußballfeldgroßer Saal im Rechenzentrum des LDS NRW zeugt noch heute von diesen monumentalen Anlagen. Die Anwender hingegen haben an Terminals gearbeitet, die sternförmig an die Großrechner angebunden waren. Sie wurden als Eingabe- und Datensichtgerät verwendet, während die gesamte Rechenlogik auf den Großrechnern ausgeführt wurde. Die Terminals beschränkten sich dabei auf eine rein alphanumerische Ein- und Ausgabe. Die Bandbreite der Netze war zu dieser Zeit gering, allerdings waren auch die Anforderungen an die Software und die alphanumerische Ein-/Ausgabe entsprechend niedrig.

Ende der 1980er-Jahre änderte sich die IT-Landschaft grundlegend. Es kam die Zeit der grafischen Betriebssysteme. Neue Softwareanwendungen wurden für die Nutzer einfach und komfortabel bedienbar. Damit stieg der Bedarf an Rechenleistung allerdings gewaltig an. Der Siegeszug des Personalcomputers stand bevor. In der IT-Infrastruktur gab es einen Paradigmenwechsel: Die Anwendungen wurden nicht mehr zentral auf dem Großrechner, sondern lokal auf dem Arbeitsplatz-PC ausgeführt. Die Rechenleistung wurde also dort bereitgestellt, wo sie tatsächlich benötigt wurde, nämlich am einzelnen Arbeitsplatz der Kollegin bzw. des Kollegen. Parallel dazu gab es auch im Netzwerkbereich bedeutende Entwicklungen. Die Netze wurden leistungsfähiger und stellten mehr Bandbreite zur Verfügung. Durch strukturierte Verkabelung, aktive Netzwerkkomponenten und intelligente Netzwerkprotokolle stieg die Leistungsfähigkeit der Netze enorm an. Die Entwicklung gipfelte in dem Siegeszug des Ethernet-Standards und des TCP/IP-Protokolls. Wichtige Funktionen wurden nun innerhalb der Netze zentral bereitgestellt. Auf leistungsfähigen Servern wurden zentrale Dienste, wie beispielsweise Datenhaltung (Fileserver), Druckdienste (Printserver), Anwendungen (Applikationsserver) und Datenbanken (Datenbankserver), zur Verfügung gestellt. Dies beschreibt das bis heute am weitesten verbreitete sog. „Client-Server-Computing“. Auf unseren zahlreichen Arbeitsplatz-PCs werden Office-Anwendungen, Datenbankanwendungen etc. ausgeführt. Die Daten werden vielerorts auf zentralen Fileservern abgelegt, über einen Printserver werden Druckdaten an den Netzwerkdrucker geschickt. Viele zentrale Server arbeiten aber auch im Hintergrund und ihre Nutzung ist dem Anwender oft nicht bewusst. Wir melden uns morgens an einer Windows-Domäne (Domänencontroller) an, gehen über einen Proxyserver ins Internet und fragen Internet-Adressen über einen Nameserver (DNS-Server) ab. Somit gibt es eine Vielzahl von Diensten, die zentral im Netzwerk bereitstehen und im Hintergrund ablaufen.

Heutige Anforderungen

Was kennzeichnet den Umgang mit unserer heutigen IT-Infrastruktur? Was sind Leistungsmerkmale und wo stoßen wir an Grenzen?

Die IT-Infrastruktur ist heute sehr vielfältig, heterogen und vor allem eines: Sie ist enorm komplex. Die Anzahl der Server und der bereitgestellten Dienste in den Netzen ist stark angewachsen. Die Anzahl von PC-Arbeitsplätzen ist immens groß. Manche Kollegen besitzen für verschiedene Aufgaben sogar mehr als einen PC. Die Aufgabe der IT-Abteilungen ist es, diese Infrastruktur zu managen. Für die Bereitstellung und Wartung der einzelnen IT-Systeme benötigen sie dazu gut ausgebildetes und spezialisiertes Personal. Sowohl die hohe Anzahl an benötigten IT-Systemen als auch der Personalaufwand haben zu steigenden Kosten geführt. Es ist sehr aufwendig und technisch schwierig geworden, eine Infrastruktur von dieser Komplexität technisch zu handhaben. Ferner steht für den Betrieb nur ein begrenztes Budget zur Verfügung, welches nicht weiter wächst bzw. sogar sinkt. Damit stößt die derzeitige IT-Landschaft an Grenzen. Es wird der Ruf nach einer Umgestaltung laut. Die IT-Infrastruktur muss einfacher, übersichtlicher und leichter zu handhaben sein. Und damit sind wir wieder am Anfang, bei dem Begriff der Konsolidierung.

In der Landesverwaltung NRW findet sich eine Vielzahl von IT-Anwendungen. Office-Anwendungen auf PCs dienen den regulären Geschäftsabläufen innerhalb unserer Behörden und Einrichtungen. Daneben gibt es vielfältige IT-Fachanwendungen, die innerhalb einer Behörde, behördenübergreifend, ressortübergreifend oder sogar landesweit zum Einsatz kommen. Hinzu kommt eine IT-Infrastruktur, die häufig auf mehrere Standorte verteilt ist. Eine Konsolidierung dieser IT-Landschaft bedeutet, IT-Systeme zu zentralisieren, zu standardisieren, sie zu vereinfachen, den Administrationsaufwand zu verringern und damit Kosten einzusparen. Dieser Trend ist bereits seit einigen Jahren zu beobachten. Wie mit der Schlüsseltechnologie Server Based Computing eine praktikable Lösung genutzt werden kann, wird im Folgenden erörtert.

Server Based Computing

Server Based Computing (SBC) bedeutet die Bereitstellung von Daten und Anwendungen auf zentralen Servern oder Serverfarmen. Die Anwendungen laufen dabei nicht mehr wie gewohnt auf einem Arbeitsplatz-PC ab, sondern auf dem entsprechenden Server. Der lokale Rechner (Client) wird lediglich für die Bildschirmausgabe und die Eingabe von Tastatur- und Mauskommandos benötigt. Die Anwendungen werden auf diese Weise zentral an einem Punkt im Netz bereitgestellt.

Diese Technologie bietet eine ganze Vielzahl von Vorteilen, die wir uns vor Augen führen wollen. Als Erstes muss die Anwendung nur ein einziges Mal zentral installiert werden. Kein Administrator muss die Anwendung bei den Nutzern vor Ort einrichten. Die umständliche sog. „Turnschuh-Administration“ gehört der Vergangenheit an. Das Gleiche gilt für sämtliche Änderungen an der Anwendung. Eine Aktualisierung auf eine neue Version, eine Erweiterung oder ein Sicherheitsupdate ist in kürzester Zeit möglich, weil diese Maßnahmen einmalig zentral durchgeführt werden. Dadurch ist gewährleistet, dass alle Nutzer unmittelbar und zeitgleich mit der aktuellen Version arbeiten. Ist eine Anwendung einmal veröffentlicht und der Zugriff für die Nutzer eingerichtet, reduziert sich auch die Nutzerbetreuung (Support) auf ein Minimum. Mögliche Probleme einer Installation auf dem Client sind von vornherein ausgeschlossen. Die Wartung der Anwendung auf den Servern kann fernadministrativ (remote) durchgeführt werden. Ein weiterer Vorteil ist die Integrität der Anwendung. Sie ist in einer homogenen Umgebung installiert und muss nicht mit anderen Anwendungen auf dem Client zusammenarbeiten. Damit sind mögliche Kompatibilitätsprobleme ausgeschlossen. Die Sicherheit der Anwendung wird ebenfalls verbessert. Sie wird nicht durch offen liegende Schnittstellen und einen

möglicherweise unzureichenden Virenschutz auf dem Client beeinträchtigt. Durch eine sog. „Härtung“ der zentralen Server mittels diverser Maßnahmen (zentraler Virenschutz, Zugriffsbeschränkungen, Gruppenrichtlinien etc.) kann die Sicherheit weiter optimiert werden. Das heißt, es wird mit Server Based Computing eine leistungsfähige IT-Infrastruktur geschaffen, die für die IT-Administration und für den Anwender enorme Vorteile bietet.

Last but not least stellt sich immer auch die Frage nach den Kosten einer IT-Infrastruktur mit Server Based Computing. Die Kosten sind für die Entscheidung für oder gegen eine IT-Architektur oder ein IT-Verfahren maßgebend und werden daher in einem gesonderten Abschnitt betrachtet. Zunächst aber werden die technische Plattform mit den verwendeten Produkten und die Infrastruktur im LDS NRW vorgestellt.

Die Plattform für SBC

Der erste Bedarf an dieser Technologie entstand bereits vor einigen Jahren im Sachgebiet LAN-Server-Administration im LDS NRW. Hier werden Infrastrukturserver für Kunden und Einrichtungen des Landes bereitgestellt und betreut. Einige Beispiele zu Infrastrukturservern sind Domänencontroller, Nameserver (DNS-Server), Fileserver und E-Mail-Server (Exchange-Server). Diese Server stellen grundlegende Dienste in Netzen zur Verfügung. Ein technischer Schwerpunkt liegt dabei auf den Betriebssystemen Microsoft Windows 2000 Server und Windows Server 2003. Diese Betriebssysteme bringen bereits von Haus aus die Windows Terminal Services mit, mit denen einfache Server Based Computing-Anwendungen realisiert werden können. Hier konnte auf der hohen Kompetenz der Mitarbeiterinnen und Mitarbeiter im Bereich der Microsoft Betriebssysteme aufgebaut werden.

Erste Anwendungen wurden mit diesen Diensten realisiert. Als daraufhin der Bedarf an weiteren Anwendungen anstieg, wurde es notwendig, sie auf einer gemeinsamen Serverplattform zu integrieren. Hier wurden mit den Windows Terminal Services jedoch sehr schnell die technischen Grenzen erreicht. Hilfe brachte der Einsatz von MetaFrame des Herstellers CITRIX. MetaFrame baut auf den Windows Terminal Services auf und ergänzt diese durch umfangreiche Funktionalitäten. Damit war eine Basis für den sicheren und einfach zu administrierenden Betrieb unterschiedlicher Anwendungen auf einer gemeinsamen Serverplattform gefunden. Die Vorteile dieser Plattform werden im Folgenden ergründet. CITRIX MetaFrame baut auf dem Konzept der so genannten „Serverfarm“ auf. Dies bedeutet, dass mehrere Server die Aufgaben der zentralen Bereitstellung von Anwendungen gemeinsam übernehmen. Sie sind zu einer logischen Einheit zusammengeschlossen. Innerhalb der Serverfarm stellen mehrere Server die gleiche Anwendung bereit. Durch diese redundante Auslegung ist zugleich die wichtige Anforderung der Ausfallsicherheit erfüllt. Im Falle des Ausfalls eines einzelnen Servers innerhalb der Serverfarm kann ein anderer Server die sog. „Anwendungssitzungen“ (Sessions) der Nutzer übernehmen. Ist die Gesamtkapazität der Serverfarm hinreichend dimensioniert, steht einem sicheren Betrieb der Anwendung nichts mehr im Wege. Der Einsatz einer Anwendung auf mehreren Servern ermöglicht zudem eine Lastverteilung (Loadbalancing). Der Zugriff eines Nutzers auf die Serverfarm wird in diesem Fall an den Server mit der geringsten Auslastung weitergegeben. Somit stehen alle Server unter der gleichen Last und die Kapazität der Serverfarm kann effizient verwaltet werden (Skalierbarkeit). Würde die Serverfarm mit ihrer Kapazität an Grenzen stoßen, kann ein weiterer Server mit den entsprechenden Anwendungen in die Farm aufgenommen werden. Damit ist eine schnelle Erweiterbarkeit

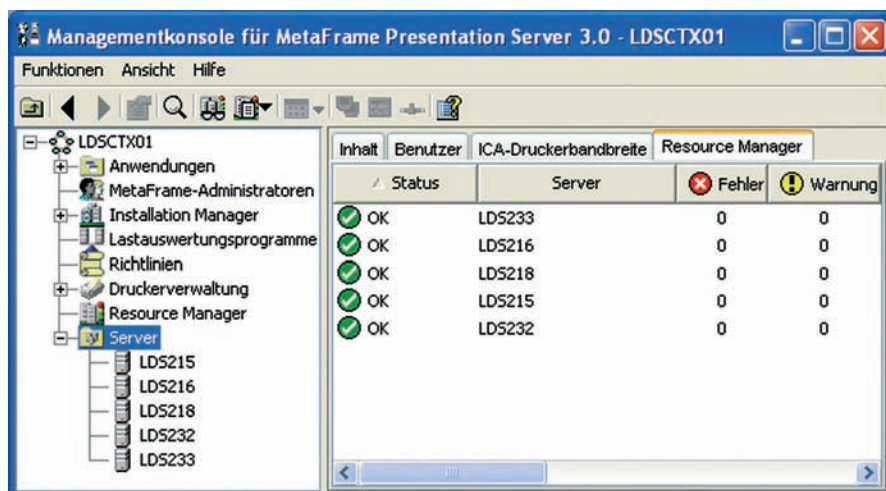


Abb. 1: Die CITRIX Management Console

und ein kontrolliertes Wachstum der Serverfarm gegeben. CITRIX MetaFrame bringt leistungsfähige Administrations- und Überwachungswerkzeuge (Tools) mit, um diese und andere Funktionen der Serverfarm zu verwalten und zu optimieren. Mit der „CITRIX Management Console“ (Abb. 1), dem zentralen Verwaltungstool von MetaFrame, sind die Administratoren stets zeitaktuell über den Zustand der Serverfarm, der Server und der freigegebenen Anwendungen informiert. Mit Hilfe der Management Console können Anwendungen freigegeben, einzelnen Nutzern zugewiesen und deren Betrieb überwacht werden. Die Management Console bringt aber auch die Möglichkeit mit, Anwendungen automatisch zu installieren und sofort auf mehrere Server der Farm gleichzeitig zu verteilen. Insbesondere in einer großen Serverfarm lässt sich so der Verwaltungsaufwand minimieren. Eine weitere Funktionalität ist die Spiegelung von Nutzersitzungen (Session shadowing). Diese erlaubt es unseren Administratoren, sich auf die Session eines Nutzers aufzuschalten, um Probleme gemeinsam mit dem Nutzer direkt an der angezeigten Programmoberfläche zu beheben. Die Behebung von Fehlern (Troubleshooting) wird damit wesentlich erleichtert. Alles in allem bietet uns die Management Console damit ein umfangreiches Set von Funktionalitäten, das für den Betrieb der Serverfarm kaum Wünsche offen lässt.

Als Dienstleister der Landesverwaltung ist das LDS NRW bestrebt, seinen Kunden innovative und technisch optimale Lösungen anzubieten. Mit dem Konzept der Serverfarm kann dem Rechnung getragen werden. Die Serverfarm ermöglicht die Bereitstellung mehrerer Anwendungen auf einer gemeinsamen Plattform, ohne dass sich die Anwendungen im laufenden Betrieb gegenseitig beeinträchtigen. Damit ist die Möglichkeit gegeben, IT-Verfahren und Anwendungen für mehrere Kunden auf einer gemeinsamen Plattform effizient und damit kostengünstig bereitzustellen.

Die CITRIX-Infrastruktur

Nachdem nun die Grundlagen der technischen Plattform bekannt sind, wird die CITRIX-Infrastruktur näher betrachtet, so wie sie im Rechenzentrum des LDS NRW betrieben wird (Abb. 2).

Den ersten Baustein bildet der Domänencontroller. Er ist für die logische Organisation der Server in einer Domäne des Active Directory nrw.de verantwortlich. Alle Server der CITRIX-Infrastruktur sind Windows-Server mit dem Betriebssystem MS Windows 2000 oder Windows Server 2003 und sind in der Domäne „service.lds.nrw.de“ positioniert. Damit im Fehlerfall nicht die ganze Domäne zum Erliegen kommt,

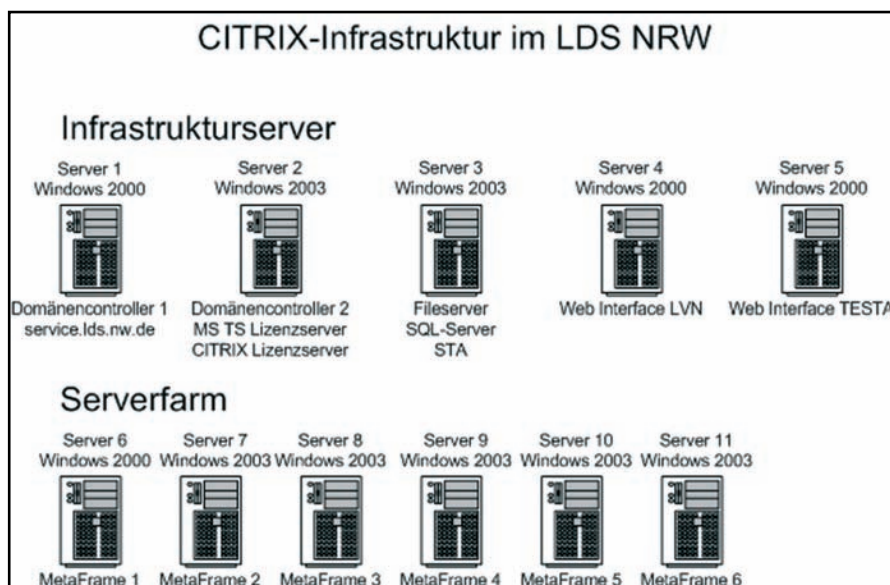


Abb. 2: Die CITRIX-Infrastruktur

wurde die CITRIX-Infrastruktur mit einem zweiten Domänencontroller ausgestattet.

Daneben wird ein Fileserver eingesetzt. Er verwaltet den zentralen Datenbestand (Datastore) der Serverfarm bzw. der Serverfarmen. Der Datastore ist eine Datenbank, in der alle relevanten Daten über die Serverfarmen abgelegt sind. Hier finden sich auch die Informationen wieder, die über die CITRIX Management Console eingegeben und abgerufen werden. Gleichzeitig werden auf dem Fileserver die Daten der Nutzer der Anwendungen verwaltet; die Benutzerkonten (Profiles) aller Nutzer der Serverfarmen sind hier registriert.

Des Weiteren befinden sich zwei Web-Server im Einsatz. Hintergrund für die Nutzung dieser Server ist die Problemstellung, dass sich auf dem Verbindungsweg durch das LVN zwischen Client und Serverfarm häufig Firewalls für die IT-Sicherheit befinden. CITRIX nutzt für die Datenkommunikation zwischen Client und Serverfarm das CITRIX-eigene so genannte ICA-Protokoll (Independent Computing Architecture). Dieses Protokoll verwendet den TCP/IP-Port 1494, der auf Firewalls häufig gesperrt ist. Um hier den Verwaltungsaufwand zu minimieren und diese technische Hürde zu überwinden, sind zwei Produkte von

CITRIX im Einsatz, nämlich das CITRIX Web Interface und das CITRIX Secure Gateway. Das Web Interface bietet die Möglichkeit, den Datenverkehr nicht über das ICA-Protokoll, sondern über das Protokoll http (TCP/IP-Port 80) zu führen. Der Aufruf der Anwendung erfolgt durch den Nutzer aus einem Webbrowser heraus, was weitere Vorteile bietet. Die Hürde der Firewalls kann so überwunden werden, da die Firewalls im Regelfall für den Datenverkehr über den TCP/IP-Port 80 freigeschaltet sind. Das Secure Gateway bietet darüber hinaus die Möglichkeit, den Datenverkehr mit 128 bit zu verschlüsseln und über das Protokoll https (TCP/IP-Port 443) zu führen. Dieser steht ebenfalls über Firewalls hinweg zur Verfügung. Das Web Interface und das Secure Gateway lassen sich auf einem gemeinsamen Server integrieren. Allerdings mussten aus tech-

nischen Gründen zwei Server eingesetzt werden; der eine für den Datenverkehr im Landesverwaltungsnetz (LVN) und der andere für das TESTA-Netz (Trans-European Services for Telematics between Administrations). Letzteres wird für den Datenaustausch mit Kommunen und anderen Bundesländern benötigt. Beide Server stellen jeweils ein Web Interface und ein Secure Gateway bereit. Sie beinhalten jeweils ein SSL-Server-Zertifikat, welches für den Aufbau der verschlüsselten Verbindungen genutzt wird.

Um CITRIX-Serverfarmen zu betreiben, ist noch eine weitere Komponente erforderlich, der Lizenzserver (Abb. 3). Er ist für die Verwaltung und Zuweisung der erworbenen und eingesetzten Lizenzen zuständig. Hierbei ist zwischen Lizenzen für den Microsoft Terminalserver und für CITRIX MetaFrame zu unterscheiden. Greift ein Nutzer auf die Serverfarm zu, ist zunächst eine MS Terminalserver-Verbindungslicenz (Terminal Server Client Access License – TS-CAL) erforderlich. Diese wird einmalig an den Client vergeben und steht dann für den Betrieb zur Verfügung. Für den Zugriff auf MetaFrame wird dann noch eine MetaFrame-Lizenz benötigt. Diese Lizenz wird nur während des Zugriffs auf die Serverfarm vergeben und wird nach dem Abmelden wieder frei. Es handelt sich um Lizenzen für gleichzeitige Nutzer (Concurrent user). Hier unterscheidet sich die Lizenzpolitik von Microsoft und CITRIX. Dies führt dazu, dass die Menge der benötigten CITRIX-Lizen-

Produkt	Modell	Typ	Installiert	In Verwendung	Verfügbar	% in Verwendung
Citrix Startlizenz	Server	System	5.000	11	4.989	0,2
MetaFrame Presentation Server, Enterprise Edition	Gleichzeitige Benutzer	Retail	205	16	189	7,8

Abb. 3: Die CITRIX License Management Console

zen eine Teilmenge der TS-CALs bilden. Die Standardfrage lautet: Wie viele der potenziellen Nutzer greifen gleichzeitig auf die Serverfarm zu? Wie später noch gezeigt wird, ist diese Frage bei der Ermittlung der Kosten für den Betrieb einer Anwendung von Bedeutung. Beide Lizenzserver laufen als Serverdienst auf dem zweiten Domänencontroller.

Neben den Infrastrukturservern werden die MetaFrame-Server eingesetzt, die zusammen die Serverfarm bilden. Sie sind das Herzstück der technischen Plattform. Hier sind die Anwendungen der LDS-Kunden installiert. Die Server müssen sehr leistungsfähig sein. Sie sind daher großzügig dimensioniert und auf eine besonders hohe Performance ausgelegt. Als Konfigurationsbeispiel sei ein Server des Typs HP DL 380 mit 2 Intel 3.0 GHz Xeon-Prozessoren und 4 GB Arbeitsspeicher genannt. Als Softwarestandard wird derzeit Windows Server 2003 und CITRIX MetaFrame Presentation Server 3.0 eingesetzt. Mit dieser Konfiguration ist sichergestellt, dass die Server viele Sessions gleichzeitig verarbeiten können. Die Anzahl der von einem Server zu verarbeitenden Sessions ist für die Dimensionierung der Server bzw. der Serverfarm entscheidend.

Letztlich ist diese Zahl auch für die Ermittlung der Hardwarekosten maßgebend. CITRIX nennt einen Richtwert von 50 Sessions pro Server. Da Anwendungen in erforderlichem Arbeitsspeicher und benötigter Rechenleistung jedoch höchst unterschiedlich sind, kann sich letztlich nur in einer Testumgebung oder im realen Betrieb zeigen, wie viele Sessions ein Server tatsächlich verarbeiten kann.

Der Citrix-Client bei den Nutzern

Wie sieht die Anwendung auf meinem Arbeitsplatz-PC aus? Kann ich komfortabel mit einer Anwendung arbeiten, die fernab meines Büros in einem Re-

chenzentrum läuft? Dies sind häufig an uns gestellte Fragen, die wir im Folgenden beantworten wollen.

Für den Nutzer bietet CITRIX drei verschiedene Client-Programme an, die den Zugriff auf die veröffentlichten Anwendungen einer Serverfarm ermöglichen. Alle Clients bieten Zugriff auf die Anwendungen, unterscheiden sich jedoch in ihrem Funktionsumfang und den genutzten Netzwerkprotokollen.

Der Standardclient ist der „Program Neighborhood“ (auch ICA-Client). Er bietet den größtmöglichen Funktionsumfang und ist daher für versierte Nutzer geeignet. Für den Zugriff auf die Serverfarm verwendet er das ICA-Protokoll (TCP/IP-Port 1494). Daneben gibt es den Web-Client (Abb. 4). Der Web-Client kann gegenüber dem ICA-Client von einem Internet-Browser aus aufgerufen werden, wobei der MS Internet Explorer oder der Netscape Browser verwendet werden kann.

Beim erstmaligen Aufruf der Anwendung im Browser wird der Web-Client im Hintergrund installiert. Die Verwendung des Web-Clients erfordert auf der Serverseite den Einsatz des CITRIX Web Interface. Hier wird der Web-Client bereitgestellt und kann somit auf einfache Art und Weise verteilt werden. Im Unterschied zum ICA-Client bietet der Web-Client weniger Anpassungs- und Einstellungsmöglichkeiten, ist daher aber für einen größeren Nutzerkreis geeignet. Da der Web-Client die breitesten Einsatzmöglichkeiten bietet und das LDS NRW das Web Interface einsetzt, ist der Web-Client der bevorzugte Client. Als kleines Tool bietet CITRIX noch den „Program Neighborhood Agent“ (Abb. 5) an. Er integriert den Zugriff auf die Anwendungen in den Windows-Desktops. Mit Hilfe von Icons lassen sich die Anwendungen komfortabel aufrufen. Der Program Neighborhood Agent benötigt allerdings das Web Interface im Hintergrund. Es besteht aber

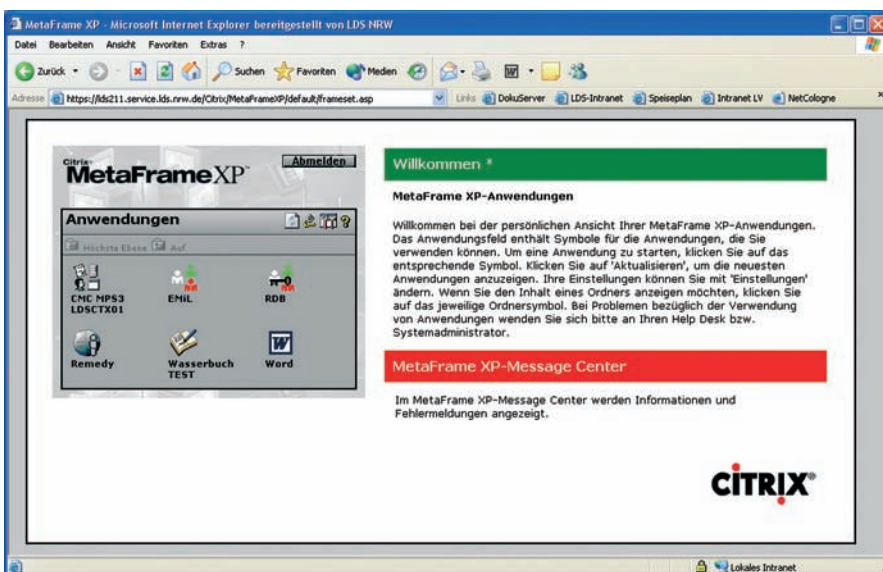


Abb. 4: Der CITRIX Web-Client

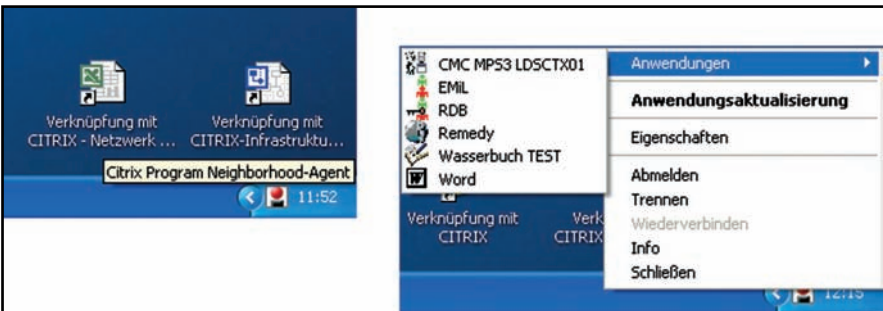


Abb. 5: Der CITRIX Program Neighborhood Agent

auch die Möglichkeit, mit einem Client für Java auf die Anwendungen zuzugreifen. Dazu wird ein Internet-Browser mit den Erweiterungen Java von Sun Microsystems oder der Microsoft Virtual Machine benötigt. Dieser Zugang hat jedoch einen geringen Funktionsumfang und bietet sich nicht für einen professionellen Einsatz an. Er ist vorwiegend für IT-Umgebungen vorgesehen, in denen vielfältige Client-Betriebssysteme eingesetzt werden. In einer überwiegend homogenen Microsoft Infrastruktur macht der Einsatz des Clients für Java wenig Sinn. Allerdings unterstützt er auch alternative Internet-Browser, wie beispielsweise den Mozilla Firefox.

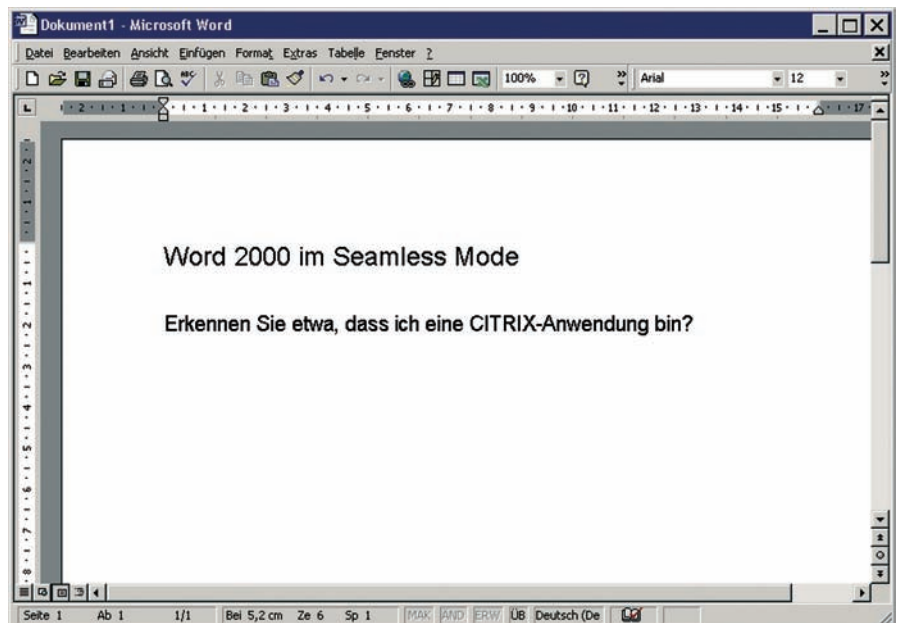


Abb. 6: MS Word 2000 unter CITRIX MetaFrame

Bei der Betrachtung der Client-Infrastruktur darf die Netzwerktechnik nicht unerwähnt bleiben. Leistungsfähige Netze mit hoher Verfügbarkeit haben das Server Based Computing überhaupt erst ermöglicht. Der reguläre Betrieb einer CITRIX-Session ist dabei nicht sehr ressourcenbindend. Im regulären Betrieb werden lediglich Bildinformationen sowie Tastatur- und Mauseingaben verarbeitet. Daher benötigt eine Session im Regelfall nur 20Kbit/s an Bandbreite. Schon mit einer einfachen ISDN-Leitung mit 64Kbit/s-Datenübertragungsrate lässt sich komfortabel arbeiten. Über den regulären Betrieb hinaus geht allerdings die Verarbeitung von Druckdaten. Wenn ein Dokument aus einer CITRIX-Session heraus auf einem lokalen Arbeitsplatzdrucker ausgedruckt werden soll, fließt kurzzeitig eine hohe Datenmenge über die Leitung. Doch auch für solche Sonderfälle bietet CITRIX Lösungen an. Mit besonderen Kompressionsverfahren wird die Datenmenge komprimiert und die Leitung weniger belastet.

Um einen Eindruck von dem Aussehen einer CITRIX-Session zu gewinnen, wird im Folgenden als Beispiel die Anwendung MS Word 2000 betrachtet (Abb. 6). Im besten Fall fällt es dem Nutzer gar nicht auf, dass er

auf einer Serverfarm arbeitet. Die Anwendung sollte so aussehen, als wenn sie auf seinem lokalen Arbeitsplatz-PC abläuft. Hierzu bietet CITRIX den „Seamless Mode“ an. Mit dieser Funktion wird eine CITRIX-Session lokal in einem Fenster ausgeführt, dessen Größe wie bei lokalen Anwendungen verändert werden kann. Das Sitzungsfenster kann verschoben, minimiert und maximiert werden. Darüber hinaus kann Text aus einer lokalen Anwendung in die CITRIX-Anwendung und umgekehrt kopiert werden. Alles in allem wird somit eine Funktionalität geboten, die einer lokalen Anwendung voll entspricht und zusätzlich die Vorteile einer zentralen Steuerung bietet, wie sie im vorigen Abschnitt bereits erläutert wurden.

Anwendungsbeispiele

Wo bieten sich in Ihrer Behörde bzw. in ihrem Betrieb Einsatzmöglichkeiten für Server Based Computing? Wie können Sie den optimalen Nutzen aus dieser Technologie ziehen? Die Möglichkeiten sind in der Tat vielfältig. Vieles hängt von der gegebenen IT-Infrastruktur und der Organisation der bestehenden IT-Verfahren ab. Im Fol-

genden sollen drei Anwendungsbeispiele aufgezeigt werden.

Beispiel 1: Sie betreiben eine Datenbank-anwendung in mehreren Lokationen. Dabei müssen Sie mehrere Datenbanken auf Datenbankservern an verschiedenen Standorten auf dem gleichen Stand halten. Die Installation der Datenbankclients auf den lokalen Arbeitsplatz-PCs der Nutzer ist zudem aufwändig und zeitintensiv. Ein Rollout einer neuen Version des Datenbankclients ist ohne zusätzlichen Personaleinsatz nicht möglich.

Nichts liegt hier näher, als über eine Zentralisierung und Konsolidierung hin zu einer Datenbank in einem zentralen Rechenzentrum nachzudenken. Der Datenbankclient kann den Nutzern auf einer CITRIX-Serverfarm bereitgestellt werden. Die Aktualisierung dieses Datenbankclients auf eine neue Version ist nur noch eine Sache von Minuten. Die Replikation der Daten auf den Datenbankservern entfällt völlig. Sie erhalten mit nur einem Datenbankserver und mehreren CITRIX-Servern eine zentrale IT-Infrastruktur, die rationell und effizient betrieben werden kann.

Beispiel 2: Sie betreiben Arbeitsplatz-PCs mit besonders hohen Sicherheits-

anforderungen, wie zum Beispiel Arbeitsplätze von Systemadministratoren. Diese sind mit weit gehenden Benutzerrechten (Administratorrechten) für die eigenen PCs und jene der zu verwaltenden Infrastruktur ausgestattet. Um einen direkten Kontakt dieser Rechner mit dem Internet und damit einen potenziellen Angriff auf diese sensiblen Rechner zu vermeiden, können Sie Ihren Mitarbeitern den Internetzugriff von diesen Arbeitsplätzen aus nicht erlauben.

Hier ist es möglich, den Internetzugang über eine CITRIX-Serverfarm zu realisieren. Die Farm bietet dann einen zentralen Zugangspunkt zum Internet, der mit Hilfe eines Virenschanners und weiterer Schutzmaßnahmen besonders abgesichert werden kann. Im Falle eines Angriffs wird nur der Server einer Farm befallen, eine Übertragung auf die lokalen Arbeitsplatz-PCs ist nicht möglich. Mittels einer automatischen Softwareinstallation kann dieser Server im Bedarfsfall neu installiert werden und der Betrieb ist in kürzester Zeit wieder hergestellt.

Beispiel 3: Ihre Institution besteht aus einer Hauptstelle und mehreren, mitunter sehr kleinen Außenstellen, die räumlich weit voneinander entfernt sind. Sie sind aus Kostengründen mit Datenleitungen geringer Bandbreite (z. B. ISDN mit 64Kbit/s) an die Hauptstelle angebunden. Die Mitarbeiterinnen und Mitarbeiter benötigen Zugriff auf ein oder wenige Fachverfahren, die auf zentralen Servern im Rechenzentrum der Hauptstelle laufen. Der Zugriff ist häufig langsam, führt zu hohen Wartezeiten und die Beschäftigten können nicht komfortabel mit ihren Anwendungen arbeiten.

Da die Verfahren ohnehin bereits zentral an der Hauptstelle bereitgestellt werden, bietet sich hier auch die zentrale Bereitstellung der Client-Anwendung auf einer CITRIX-Serverfarm an. Die Datenmenge, die zwischen den Nutzern in der Außenstelle und den Servern in der Hauptstelle übertragen

werden muss, kann damit erheblich reduziert werden. Die Anwender können schneller mit der Datenbankanwendung arbeiten. Auf der Serverfarm kann die Anwendung komfortabel gewartet werden. Notwendige Updates müssen nur einmalig auf der Serverfarm durchgeführt werden. Kein Beschäftigter muss in eine Außenstelle fahren, um dort die Datenbankanwendung auf dem Arbeitsplatz-PC eines Nutzers zu aktualisieren. Damit steht dieses Personal für andere Aufgaben zur Verfügung.

Wie Sie sehen, ist der Einsatz von Server Based Computing mit CITRIX MetaFrame immer dort sinnvoll, wo umfangreiche Konsolidierungsmaßnahmen möglich sind. Die Komplexität der Infrastruktur wird reduziert und der Betrieb wird kostengünstiger. Die neue Infrastruktur kann mit geringerem Personalaufwand betrieben werden.

Die Kosten

Die Entscheidung für oder gegen ein IT-Projekt ist in allererster Linie eine Frage der Kosten. Die Budgets für den Betrieb und die Weiterentwicklung der IT-Infrastruktur sind begrenzt. Daher müssen neue Investitionen besonders vor dem Hintergrund betrachtet werden, inwieweit sie sich amortisieren und langfristig zu einer Kosteneinsparung führen. Das Server Based Computing hat sich in den letzten Jahren weit verbreitet. Dies liegt vor allem daran, dass mit dieser Technologie langfristig hohe Einsparungen erzielt werden können. Wie ist das möglich?

Die Kosten einer IT-Infrastruktur bestehen zum einen aus den Anschaffungskosten für die Hardware (Server, Netzwerk, PCs, Wartungsverträge etc.) und die Software (Lizenzen für Betriebssysteme und Anwendungssoftware, Pflegeverträge etc.). Zum anderen entstehen Kosten für den Betrieb der IT-Infrastruktur (Systemmanagement,

Wartungsarbeiten, Fehlerdiagnose, Nutzerbetreuung etc.). Hierbei zeigt sich, dass die Kosten für den Betrieb der herkömmlichen IT-Infrastruktur die Kosten für die Anschaffung um ein Vielfaches übersteigen. Für den Betrieb der IT-Infrastruktur ist ein hoher Personalaufwand erforderlich, welcher einen Großteil der Kosten ausmacht. An dieser Stelle erzielt der Einsatz von Server Based Computing die größten Einspareffekte. Durch die Bereitstellung der Anwendungen in einem zentralen Rechenzentrum entfällt die aufwändige Installation und Unterstützung an den Arbeitsplätzen der Nutzer. Der gesamte Betrieb wird von einer zentralen Stelle aus durchgeführt, und das ist höchst effizient. Server Based Computing kann zudem über die bestehende IT-Infrastruktur hinweg eingesetzt werden. Die Arbeitsplatz-PCs der Nutzer werden weniger belastet. Damit sinken die Innovationszyklen der PCs, sie müssen ggf. nicht so schnell durch neue Geräte ersetzt werden. Die Netze werden ebenfalls entlastet. An diesen Stellen werden also weitere Kosteneinsparungen erzielt. Dem stehen die Anschaffungskosten für die zentrale Serverfarm und die Kosten für CITRIX- und Terminalserver-Lizenzen gegenüber. Alles in allem hängt es jedoch von den konkreten Ausgangsbedingungen eines IT-Verfahrens ab, wie hoch die zu erzielenden Einsparungen im Einzelnen sind und ob sich der Einsatz lohnt. Schätzungen im Bereich der TCO-Analyse (Total Cost of Ownership) sprechen dabei von Einsparpotenzialen bis zu einer Höhe von 60 Prozent.

Unsere Leistungen für die Landesverwaltung

Das LDS NRW hat eine leistungsfähige Infrastruktur für den Betrieb von Anwendungen für Kunden der Landesverwaltung aufgebaut. Unser Ziel ist es, diese optimal in allen Bereichen des Einsatzes von Server Based Computing

zu unterstützen. Die Dienstleistungen reichen von Beratung zum Einsatz von Server Based Computing über den Test von Anwendungen auf ihre Eignung für die zentrale Bereitstellung bis hin zum Betrieb der Anwendungen auf unserer CITRIX-Serverfarm.

Hierzu stehen leistungsfähige Mitarbeiterinnen und Mitarbeiter mit einer fundierten Ausbildung zur Verfügung, die ihre Kompetenz durch die regelmäßige Teilnahme an Veranstaltungen bei Herstellern, Lieferanten und Fortbildungsdienstleistern auch zukünftig sicherstellen. Und vor allem können wir

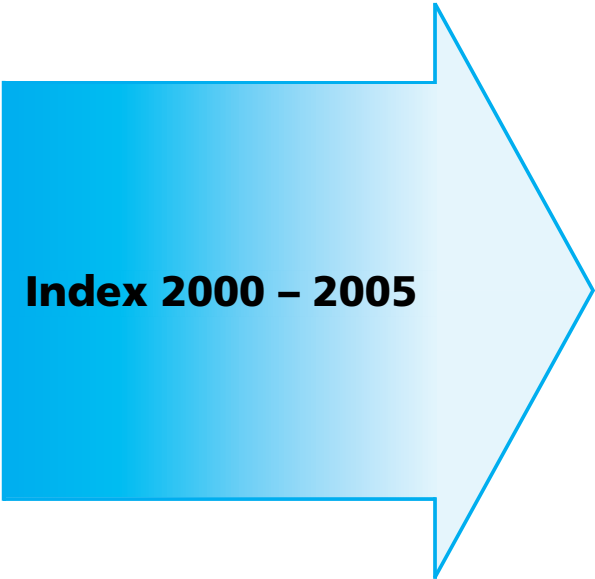
auf die mittlerweile mehrjährige Erfahrung im Aufbau und Betrieb der CITRIX-Infrastruktur im LDS NRW zurückblicken.

Das LDS NRW bietet an, mit Ihnen gemeinsam die Möglichkeiten eines Einsatzes von Server Based Computing für Ihre IT-Anwendungen zu erörtern und zu prüfen. Hierzu wird eine detaillierte Analyse der Rahmenbedingungen Ihrer IT-Infrastruktur durchgeführt. Ein Testverfahren kann zeigen, ob Ihre Anwendungen für die zentrale Bereitstellung geeignet sind. Eine Kosten-Nutzen-Analyse, eine Gegen-

überstellung der Vor- und Nachteile und ein Dienstleistungsangebot sollen die Grundlage für Ihre Entscheidung sein, ob das LDS NRW als Dienstleister für Sie tätig werden kann. Begeben Sie sich mit Server Based Computing auf einen sicheren Weg in die Zukunft der IT-Infrastruktur!



*Dipl.-Wi.-Ing.
Sascha Bittau
Tel.: 0211 9449-3585
E-Mail: sascha.bittau
@lds.nrw.de*



Index 2000 – 2005

2/2005	Wahlen in NRW – Die Unterstützung der Landeswahlleiterin durch das LDS NRW Dr. Thomas Pricking
	Landtagswahl 2005 – Interaktive Kartendarstellung der Wahlergebnisse Anja Neumann, Georg Stahl, Christoph Rath
	Aktuelle Trends in der Spracherkennung Dieter Bittner
	JUDICA/TSJ – Neue Zusammenarbeit mit dem Justizministerium Kirsten Rölleke, Claudia Schröder
	Aufwandsschätzungen in Softwareprojekten Dr. Jan Mütter, Ulrich von Hagen
	Risikomanagement in Softwareprojekten Ulrich Andree, Dr. Jan Mütter
	Projektmanagement in der Praxis oder: Wo steht mein Projekt? Thomas Reiff (IBM), Dr. Jan Mütter
	Server Based Computing – Auf dem Weg in die Zukunft Sascha Bittau

1/2005	Einsatz von Open Source Software im LDS NRW Dr. Frank Dillmann, Susanne Schmitz, Dr. Thomas List
	Web Service – Technologie der Zukunft? Dr. Heike Wellmeyer
	Serverbereitstellung in der Landesdatenverarbeitungszentrale Dr. Michael Fluchtmann
	Einrichtung sicherer Verzeichnisse für Arbeitsgruppen mit Windows-Technologien Stefan Willer, Holger Traczinski, Dr. Jörg Hintelmann
	Mono: .Net für Linux Dr. Susanne Wigard
	WAVE – Eine umfassende Unterstützung für die Veranstaltungsplanung und -durchführung Monika Mika, Jürgen Sperling, Ulrich von Hagen
	Das neue V-Modell XT Ulrich von Hagen
	Agil und extrem (2) – Praktische Erfahrungen mit agilen Methoden Dr. Jan Mütter, Ulrich von Hagen

2/2004	Management moderner IT-Infrastrukturen Dr. Dirk Weckendrup, Thorsten Lentes
	Barrierefreie Webinformationsangebote Dr. Thomas Ott
	Bibliotheksverbund der Landesbehörden Anita Gruschka
	OBELIX: Die Migration der Daten Jens Andexer, Hermann Kinder
	IT-Lehrveranstaltungen – einmal anders Ulrich von Hagen, Bernhard Kruse
	Agil und extrem – neue Wege in der Softwareentwicklung Ulrich von Hagen

Ausgabe	Schwerpunktthema
1/2004	SPAM – und (k)ein Ende? Gerold Vannahme Open Source Software auf Webservern im LDS NRW Susanne Schmitz Gefahren durch Viren Kai Hesse OBELIX und die Systemtechnik Guido Winkler, Eberhard Baumgarten (IBM)
2/2003	SAPOS® – ein modernes Werkzeug zur Positionsbestimmung und Vermessung – realisiert durch Kooperation der Landesbetriebe Christian Elsner (LvermA NRW), Dr. Frank Laicher SPAM – Wenn der Postmann zwei(tausend)mal klingelt Gerold Vannahme Zentrale Problemkoordinierung durch das IT-Managementzentrum Dr. Dirk Weckendrup, Thorsten Lentes ArcGIS-Entwicklungen des Graphikzentrums im LDS NRW – Automatisch erzeugte Kartenrahmen für ArcGIS 8.x – Wolfgang Ruschke Internetzugang der Landesverwaltung NRW Dr. Christina Mendorf Landesdatenbank im Web – Status und Planung Peter Müller Datendrehscheibe – Einleiterüberwachung – Abwasser D-E-A: Stand und Entwicklungsperspektiven Dr. Eckhart Treunert, Dr. Heike Wellmeyer DV-technisches Systemdesign im Projekt Obelix Mechthild Kroll Wie hältst Du's mit den EJBs? Webanwendungen mit Java im Überblick Dr. Susanne Wigard
1/2003	GeoServer der Landesverwaltung NRW nimmt Produktion auf Stefan Küpper Angebot zentraler GIS-Dienste mit dem GeoServer Christoph Rath Landtagsinformationssystem NRW Brigitte Corves Microsofts .Net-Technologie – wohin geht die Programmierung? Dr. Susanne Wigard Programmierung mit dem VisualAge Generator im Projekt OBELIX Thomas Fischell, Jochen Gottelt Webhosting-Dienstleistungen des LDS NRW Susanne Schmitz

1/2002

**Einführung des Geoinformationssystems „ArcGIS“
in der Landesverwaltung NRW mit Unterstützung
des Grafikzentrums im LDS NRW**

Stefan Küpper

Testmanagement und Testautomatisierung im Projekt Obelix

Peter Kürschner, Ralf Gerson

Contentmanagement in NRW

**Ausbau der Webinformationsangebote der Landesverwaltung
und Erleichterung des Zugangs zu Informationen**

Dr. Thomas Ott

FÜSYS

**Ein Führungsinformationssystem zum strategieorientierten
Management erfolgsrelevanter Projekte in Ministerien
und anderen großen Behörden**

Dieter Spalink (IM NRW), Frank Walter

SAN macht Daten schnell

Storage Area Network im Windows-LAN des LDS NRW

Dr. Markus Brakmann

Die Netzbasis des Landesverwaltungsnetzes

– gerüstet für die Zukunft

Dr. Frank Laicher

'Penguin meets Dinosaur'

Linux und Großrechner finden zueinander

Dr. Frank Dillmann, Jochen Gruse

2/2001

**Das Internet als neues Medium für die Erhebung
statistischer Daten**

Dr. Thomas Pricking

**Dienstleistungen und Servicekonzepte des LDS NRW
zur IT-Infrastruktur**

Dr. Bruno Vogel

**Die Erhebung der Abwasser abgabe – DV-Unterstützung
einer gesetzlichen Aufgabe**

Rolf Büttinghaus (LUA NRW), Dr. Heike Wellmeyer

GEOSERVER der Landesverwaltung NRW im Intranet und Internet

Stefan Küpper

Der Umweltdatenkatalog NRW

**Transparente Umweltinformationen für Öffentlichkeit,
Politik und Planung**

Dr. Thomas Ott

**Kostensenkung am PC-Arbeitsplatz durch Geschäftsprozessanalyse
im Desktop Management des LDS NRW**

Dr. Martina Jansen, Klaus Monse

Bibliotheksverbund der Landesbehörden NRW

Anita Gruschka

Web Based Training im Landesverwaltungsnetz

Dr. Angela Barthen

1/2001**Kommunikation ist alles – Elektronische Post**

Maria Palarz-Kupka

Landtagsinformationssystem Nordrhein-Westfalen

Brigitte Corves

Public Key Infrastruktur – Was steckt dahinter?

Dr. Guido Cassor

Normen zur Software-Ergonomie

Dr. Elmar Kalthoff

IT-Unterstützung im Support**– Nutzung eines Trouble-Ticket-Systems**

Helmut Nehrenheim

„Gut Werkzeug – halbe Arbeit“**Unterstützung der Anwendungsbereitstellung****durch Vorgehensmodelle, Methoden und Werkzeuge**

Ulrich von Hagen

Obelix – das neue Bezügeverfahren

Dr. Jan Mütter, Ralf Gerson

2/2000**LDS NRW – Internet-Dienstleister der Landesverwaltung****Nordrhein-Westfalen**

Gerhard Bendels

Web-Anwendungen auf dem IBM-Großrechner

Peter Müller

Lehrereinstellungsverfahren in Nordrhein-Westfalen

Monika Kempf

Spracherkennungssoftware – schon einsatzbereit?

Klaus Trommer

Integriertes Netz- und Systemmanagement**– Selbstzweck oder Wunderwaffe?**

Dr. Dirk Weckendrup

1/2000**Das Büro auf Reisen****Mobilität und Datenaktualität – mobile computing Mobile Telefone, Smartphones, Handhelds, PDAs, Notebooks – eine Bestandsaufnahme**

Klaus Trommer

IT-gestützte Vorgangsbearbeitung**in der Landesverwaltung Nordrhein-Westfalen**

Annette Hochstein

DV-Vorhaben D-E-A**(Datendrehscheibe, Einleiterüberwachung, Abwasser)**

Dr. Eckart Treunert (MUNLV NRW), Dr. Heike Wellmeyer

Zertifizierung des IT-Qualitätsmanagements der LDVZ**nach DIN EN ISO 9001**

Dr. Joachim Möhring

Das Windows-NT-Netz im LDS NRW

Dr. Markus Brakmann