

Ein Technologieentwurf für eine Klasse skalierender Webapplikationen

DIPLOMARBEIT

zur Erlangung des akademischen Grades

Diplom-Ingenieur

im Rahmen des Studiums

Masterstudium Software Engineering/Internet Computing

eingereicht von

Paul Panhofer

Matrikelnummer 9725248

an der

Fakultät für Informatik der Technischen Universität Wien

Betreuung

Betreuer: Ao.Univ.Prof. Dipl.-Ing. Dr.techn. Gerald Futschek, Univ. Doz.

Mitwirkung: -

Wien, 01.10.2013

(Unterschrift Verfasser)

(Unterschrift Betreuer)

„Hiermit erkläre ich, dass ich diese Arbeit selbständig verfasst habe, dass ich die verwendeten Quellen und Hilfsmittel vollständig angegeben habe und dass ich die Stellen der Arbeit – einschließlich Tabellen, Karten und Abbildungen –, die anderen Werken oder dem Internet im Wortlaut oder dem Sinn nach entnommen sind, auf jeden Fall unter Angabe der Quelle als Entlehnung kenntlich gemacht habe.“

Kirchberg am Wagram, 3rd October 2013

ABSTRAKT

Mit dem Aufkommen von **WEB 2.0** sind wir Zeuge einer dramatischen Änderung der geforderten **ANFORDERUNGEN** an WEB APPLIKATIONEN UND SOFTWARESYSTEME im Allgemeinen geworden. Wir erfahren eine neue Art von Applikationen, die einem veränderten Katalog nicht funktionaler Anforderungen gerecht werden muß. Im Zentrum der Arbeit steht der Versuch, eine im Sinne der technischen Umsetzung **SKALIERENDE ARCHITEKTUR** für eine KLASSE VON KOMPLEXEN SOFTWAREAPPLIKATIONEN zu finden. In diesem Zusammenhang ist die Differenzierung zwischen **HORIZONTAL** und **VERTIKAL SKALIERENDEN** Lösungen essentiell und zu beachten. Als Stellvertreter für diese Klasse wird ein spezifisches Projekt gewählt, eine **BIBER WEBAPPLIKATION**.

Im Zuge der Arbeit schenken wir der **DATENVERWALTUNG** und damit dem **KONSISTENZMODELL** besondere Aufmerksamkeit. Das **CAP THEOREM** [12][16] dient in diesem Zusammenhang als fundamentaler Fokuspunkt, um die komplementären Größen und deren Beziehungen im Umfeld der Architekturdiskussion zu verstehen. Der Vergleich diverser bestehender **TECHNOLOGIEN** und **PARADIGMEN** soll ein allgemeines Verständnis für das Problem und dessen Lösung schaffen. Die gefundenen Ergebnisse dienen als Basis zur Erstellung der **SOFTWAREARCHITEKTUR**. Zur Evaluierung der Ergebnisse werden 2 **PROTOTYPEN** erstellt. Die Prototypen werden im Sinne widersprüchlicher Paradigmen entworfen und gegeneinander verglichen. Das Resultat hilft, die gewünschten Ergebnisse zu etablieren.

Die **PROTOTYPEN** weisen eine grundlegende Differenzierung in ihrem **DATEN-**, **KONSISTENZ-** und **REPLIKATIONSMODELL** auf. Ein Prototyp wird den Prinzipien der **NoSQL** Ansätze folgend gestaltet, der andere nutzt die Theorie der verteilten **RELATIONALEN DATENBANKEN** für die Datenverarbeitung. **LASTENTESTS** werden eingesetzt um das Laufzeitverhalten der Prototypen zu bewerten und zu verstehen. Das Laufzeitverhalten ist in diesem Sinne als die Summe der verschiedenen umgesetzten Paradigmen für Datenkonsistenz, Replikation und Datenmodell zu beurteilen. Im Anschluß an die Lastentests werden die generierten Meßergebnisse ausgewertet und gesichtet. Die Meßergebnisse dienen als Grundlage zur Beurteilung der eingesetzten Modelle.

Wir haben ein **PERFORMANTES LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** des SYSTEMS als primäre **NICHT FUNKTIONALE ANORDERUNG** an das SYSTEM gestellt. Die zentrale Frage ist ob es möglich ist ein solches System unter der Verwendung eines **RDBMS** Systems als Kern umzusetzen. Betrachten wir die Ergebnisse für die **SKALIERBARKEIT** des **MYSQL PROTOTYPEN** bemerken wir daß das System nur **HORIZONTAL SKALIERT**. Unser Anspruch besteht jedoch darin ein SYSTEM zur Verfügung zu stellen, daß **LINEAR SKALIERT**. Diesen Anspruch erfüllt aber nur der **NoSQL PROTOTYP**. Wir müssen feststellen daß eine **RDBMS** Lösung für unser Problem nicht in Frage kommt.

□

ABSTRACT

WEB 2.0 is becoming the BUZZ WORD in the **DISTRIBUTED SYSTEMS COMMUNITY**. The technological standards for WEB APPLICATIONS are continually changing resulting in a new stack of **NON FUNCTIONAL** requirements. The focus of the thesis is to find a **SCALEABLE ARCHITECTURE** for a special class of **DISTRIBUTED SOFTWARESYSTEMS**.

The processing of large amounts of data in a timely manner is key. The **CAP THEOREM** helps us to establish a **CLASSIFICATION** of **DATAPROCESSING DISTRIBUTED SYSTEMS**. This classification is key to understand the different **PROPERTIES** and **REQUIREMENTS** in the process of developing the **SOFTWARE ARCHITECTURE** in question.

The comparison of different classes of **DATAPROCESSING SYSTEMS** gives us an idea of the problem. **RDBMS** and **NoSQL SYSTEMS** are the current **DATAPROCESSING TECHNOLOGIES**. We have created two **SOFTWAREPROTOTYPES**. One **PROTOTYPE** has been created as a **NoSQL SYSTEM** the other one is using **RDBMS** for **DATAPROCESSING**. We have been running different tests to learn about the **PERFORMANCE**, **SCALABILITY** and **ELASTIC SPEED UP** of the **PROTOTYPES**.

We have learned that **RDBMS** as a **DATAPROCESSING TECHNOLOGY** is not able to give us the **SCALABILITY** we need. On the other hand **NoSQL** solutions are a tool that helps us to develop a **SCALABLE** architecture. □

Contents

Contents	1
List of Figures	3
1 Einführung	5
Problemstellung und Kontext der Arbeit	5
Methodisches Vorgehen	8
Technologievergleich	11
Fragestellung der Arbeit	18
Das Biber Softwaresystem - Eine Beschreibung	19
2 Anforderungsanalyse	25
Klasse von Softwareanwendungen	25
3 NoSQL: Etablierung eines neuen Datenverarbeitungs Standards	29
Einführung	29
Technologieparadigmen	31
NoSQL Datenmodelle	35
NoSQL Implementierungen	38
4 RDBMS und andere Datenbanksysteme	45
Relationale Datenbanksysteme	45
5 Implementierung der Prototypen	49
Einführung	49
Programmiersprache	50
NoSQL: Implementierung des Models	56
RDBMS - Implementierung des Models	59
6 Architekturbeschreibung	61
Eine Patternsprache	61
Top Level Architektur	63
View	68
Controller	69

Model	71
7 Durchführung der Lastentests	73
Einführung	73
Test Framework	74
Testkriterien	75
Test-Setup	76
Testdurchführung	78
Resultat	81
8 Etablierung der Ergebnisse	85
Forschungsfrage	85
Methodisches Vorgehen	87
Bibliography	91
Index	93

List of Figures

1.1	Konzeptvergleich	14
1.2	Aktorenrechte	21
1.3	Aktorenliste	22
1.4	Aktorenbeziehung	22

Einführung

PROBLEMSTELLUNG UND KONTEXT DER ARBEIT

WEB 2.0 Mit dem Aufkommen von **WEB 2.0** sind wir Zeuge einer dramatischen Änderung der geforderten Anforderungen an **WEB APPLIKATIONEN** und **SOFTWARESYSTEME** im Allgemeinen geworden. Wir sehen eine neue Art von Applikationen, die einem veränderten Katalog **NICHT FUNKTIONALER ANFORDERUNGEN** gerecht werden muß.

Viele der heute weitverbreitesten Anwendungen, finden sich in der Liste dieser Anwendungen:

- **FACEBOOK**
- **AMAZON**
- **TWITTER**
- **YOUTUBE**
- ... □

KRITISCHE NICHTFUNKTIONALE ANFORDERUNGEN Das **LAUFZEITVERHALTEN** von **CLIENT REQUESTS** wird zum entscheidenden Thema. Anwendungen sehen sich mit der Herausforderung konfrontiert Hunderttausende und Millionen von Anwendern gleichzeitig zu verarbeiten.

Das LAUFZEITVERHALTEN der Anwendung versteht sich in diesem Zusammenhang als Resultat folgender **GRÖSSEN** und **KONZEPTE**:

- **DATENVERARBEITUNG**
- **REPLIKATIONSMODELL**
- **KONSISTENZMODELL**
- **DATENSHEMA**
- **SKALIERBARKEIT** □

GRUNDLEGENDE KONZEPTE Gleichzeitig erkennen wir ein verändertes Verständnis dieser Konzepte im Lichte von **WEB 2.0**.

- **DATENVERARBEITUNG:** Im Vordergrund steht Datenverarbeitung im TERA und PETA Datenbereich. Das **SCALE UP** Potential relationaler Systeme stößt hier schnell an seine Grenzen. Erschwerend kommt hinzu daß Zugriffe nicht gleichverteilt über ein entsprechendes Zeitintervall sondern in extremen Lastspitzen auftreten.
- **KONSISTENZMODELL:** Der Stellenwert der strikten Konsistenzanforderungen und der damit verbundenen **ACID** Eigenschaften von Transaktionen tritt deutlich in den Hintergrund. Der User nimmt **EVENTUAL CONSISTENCY** in Kauf und erkaufte sich damit ein wesentlich verbessertes LAUFZEITVERHALTEN und LINEARE SKALIERBARKEIT DES SYSTEMS.

Neue Prinzipien im Datenverarbeitungsbereich werden erdacht und entwickelt um das enge Korsett der **ACID** bedingten **KONSISTENZANFORDERUNGEN** abzulegen. Die grundlegende Idee ist hier, die Restriktionen relationaler Datenbanken aufzubrechen und andere Wege in der Datenverarbeitung zu beschreiten.

- **DATENSHEMA:** In traditionellen Datenbanksystemen wird beim Einspielen von Daten ein SCHEMA erzwungen. Passen die Daten nicht zum Schema werden diese verworfen. Dieser Ansatz wird auch **SCHEMA ON-WRITE** genannt da die Daten beim Schreiben gegen das SCHEMA geprüft werden.

Verstärkt ist der Wunsch von Software Entwicklern zu bemerken die Konventionen eines **SCHEMA ON-WRITE** zugunsten eines aufgelockerten Schemaverständnisses aufzugeben. □

ALTERNATIVE ANSÄTZE Um den veränderten nichtfunktionalen Anforderungen des **WEB 2.0** gerecht zu werden müssen Lösungen gefunden werden die uns helfen uns über die gegebenen technischen Restriktionen und Paradigmen hinwegzusetzen.

Zum besseren Verständnis dieses Ansatzes sollen 2 Beispiele aus der Praxis diskutiert werden:

- **SETI@HOME:** **SETI@HOME** ist eine Anwendung die eine immense Rechenleistung aufbringt, in dem es seine Rechenlast auf beliebige Rechner im Internet verteilen kann. (siehe[20])
- **HADOOP:** **HADOOP** ist eine Framework, das in der Lage ist, komplexe Operationen auf Datenmengen beträchtlichen Umfangs durchzuführen

Trotz der Ähnlichkeit der Anforderungen an die Anwendungen, erkennen wir zwei LÖSUNGSSTRATEGIEN, die unterschiedlicher nicht sein könnten. SETI@HOME muß als Resultat seiner Softwarearchitektur wiederholt beträchtliche Datenmengen zwischen dem SERVER und den CLIENTS austauschen. Im Gegensatz dazu verfolgt HADOOP den Ansatz, den Anwendungscode zu den CLIENTS zu bringen und die Netzwerklast damit so gering wie möglich zu halten.

Im Zuge dieser Abhandlung soll ein **TECHNOLOGIEENTWURF** für eine bestimmte Klasse von Webapplikation entwickelt werden. Stellvertretend für diese Klasse von Webapplikationen wird ein System im Kontext des Wettbewerbs des **BIBERS DER INFORMATIK** entwickelt und getestet. Im Sinne eines wissenschaftlichen Diskurses wird zuerst versucht, die Anforderungen an das System zu erarbeiten. Das Verständnis der Anforderungen soll helfen, die kritischen Bereiche der Architektur zu erkennen und entsprechend zu entwerfen. Mehrere Technologien werden in diesem Zusammenhang verglichen und ihre Tauglichkeit für das System bewertet. □

METHODISCHES VORGEHEN

Zur DISKUSSION und ETABLIERUNG unserer Ergebnisse zum **TECHNOLOGIEENTWURF** des gewünschten Systems werden die folgenden Schritte durchlaufen werden:

1. DEFINITION DER GEWÜNSCHTEN KLASSE VON SOFTWAREANWENDUNGEN
2. ENTWICKLUNG DER PROTOTYPEN
3. DURCHFÜHRUNG DER LASTENTESTS
4. ETABLIERUNG UND DISKUSSION DER ERGEBNISSE □

KLASSE VON SOFTWAREANWENDUNGEN

Die für uns interessante **KLASSE VON SOFTWAREANWENDUNGEN** möchten wird durch die für sie CHARAKTERISTISCHEN EIGENSCHAFTEN beschreiben. Die gewünschte Klasse findet sich in der **DURCHSCHNITTSMENGE** der durch folgende EIGENSCHAFTEN und ANFORDERUNGEN beschriebenen Softwaresysteme:

- VERTEILTE SYSTEM
- BASIEREND AUF DEM HTTP PROTOKOLL
- UNTERSTÜTZTUNG FÜR HTML 5
- HOCH SKALIERBAR
- DATENVERARBEITUNG IM TERRABYTE BEREICH
- PERFORMANTES LAUFZEITVERHALTEN □

PROTOTYPEN

PROTOTYPEN Stellvertretend für die im Zuge der Diplomarbeit interessante Klasse von Softwareanwendungen, sollen 2 PROTOTYPEN erstellt werden. Die PROTOTYPEN werden unterschiedlichen Konzepten und Paradigmen folgend entworfen und implementiert. Ein Prototyp wird als **RDBMS** (**REALTIONAL DATABASE MANAGEMENT SYSTEM**) umgesetzt der andere verwendet **NoSQL** als DATENVERARBEITUNGSSCHICHT. Beide Lösungen müssen einem bestimmten Katalog **NICHT FUNKTIONALER** Anforderungen gerecht werden. Der Vergleich dieser unterschiedlichen Technologien und Paradigmen soll ein allgemeines Verständnis für das Problem und dessen Lösung schaffen. □

PROTOTYPENARCHITEKTUR - SCHICHTENMODELL Um eine Vergleichbarkeit der PROTOTYPEN zu gewährleisten, wird die Architektur der beiden Programme dem **SCHICHTENMODELL** folgend gestaltet. Gelingt es uns, in beiden Prototypen einen Großteil der Schichten identisch zu halten, kann ein legitimer Vergleich der Programme durchgeführt werden.

Das **SCHICHTENMODELL** etabliert eine Reihe von Prinzipien, die die Durchführung eines Vergleichs begünstigen:

- **TRENNUNG DER FUNKTIONALITÄT:** Jede Schicht implementiert unterschiedliche FUNKTIONALITÄT. Die Idee ist, die Funktionalität der einzelnen Schichten paarweise disjunkt zu halten.
- **INTERFACES:** Die Funktionalität jeder Schicht wird dem Client über ein INTERFACE zur Verfügung gestellt. Eine strikte Trennung der INTERFACES von der Implementierung hilft uns eine generische Softwarearchitektur aufzubauen.
- **ZUGRIFF:** Jede Schicht kann nur auf das Interface der benachbarten Schichten zugreifen. Dieser hierarchische Aufbau etabliert eine strikte Trennung der Funktionalität und MODULARISIERUNG des Systems.
- **MODULARISIERUNG:** Die Unterteilung der Anwendung in ein System hierarchischer Schichten etabliert gleichzeitig eine Modularisierung der Anwendung.
- **ABSTRAKTION:** Mit jeder Schicht im Schichtenmodell steigt der ABSTRAKTIONSGRAD der Software. Der durch die Abstraktion geschaffene Kontext erlaubt eine triviale Implementierung komplexer Aufgabenstellungen. □

IMPLEMENTIERUNG

MVC MODELL Das **MVC MODELL** stellt eine Interpretation des SCHICHTENMODELLS dar, eingesetzt zur Strukturierung bestimmter Klassen von Anwendungen.

Die einzelnen Schichten des Modells sind:

- **MODEL:** Das **MODEL** als Schicht ist für die DATENVERARBEITUNG der Anwendung verantwortlich
- **VIEW:** Das **VIEW** dient als ANWENDERSCHNITTSTELLE. Diese Schicht erlaubt den Zugriff des Anwenders auf die Funktionalität des Programmes.
- **CONTROLLER:** Der **CONTROLLER** dient als MEDIATOR zwischen **MODEL** und **VIEW**.

Beide Prototypen implementieren das **MVC MODELL**. Die Herausforderung in der Gestaltung beider Programme ist es sowohl das **VIEW** als auch den **CONTROLLER** für

beide identisch zu halten. Lediglich in der MODEL SCHICHT streben wir eine grundlegende Differenzierung im Sinne der Implementierung an.

Die Prototypen verwenden unterschiedliche Ansätze zur Implementierung ihrer **DATEN-**, **KONSISTENZ-** und **REPLIKATIONSMODELLE**. Ein PROTOTYP wird den Prinzipien der **NoSQL** Ansätze folgend gestaltet, der andere nutzt die Theorie der verteilten relationalen Datenbanken für die Datenverarbeitung. □

LASTENTEST

LAUFZEITVERHALTEN LASTENTESTS werden eingesetzt, um das **LAUFZEITVERHALTEN** der Prototypen zu bewerten und zu verstehen. Das LAUFZEITVERHALTEN ist in diesem Sinne als die Summe der unterschiedlichen PARADIGMEN und KONZEPTE für **DATENKONSISTENZ**, **REPLIKATION** und des **DATENMODELLS** zu beurteilen.

Beide Prototypen werden denselben LASTENTESTS unterzogen. Der Vergleich der Ergebnisse der LASTENTESTS gibt Auskunft über die Beschaffenheit des eingesetzten **TECHNOLOGIEENTWURFS**. Die gefundenen Ergebnisse helfen uns, die gewünschten Eigenschaften zu prüfen und die gewünschten Architekturentscheidung zu treffen. □

RESULTAT

Die durch die LASTENTESTS generierten Ergebnisse dienen uns als Grundlage für den gewünschten **TECHNOLOGIEENTWURF**. Zu diesem Zeitpunkt soll gleichzeitig das Ergebnis geprüft werden und der Einfluß der unterschiedlichen KONZEPTE auf das Ergebnis besprochen werden. □

TECHNOLOGIEVERGLEICH

CAP THEOREM

CAP THEOREM DAS CAP THEOREM ETABLIERT, DASS ES FÜR EIN VERTEILTES SYSTEM UNMÖGLICH IST, DIE 3 CAP EIGENSCHAFTEN (KONSISTENZ, VERFÜGBARKEIT UND ROBUST GEGENÜBER KNOTENAUSFÄLLEN) GLEICHZEITIG ZU ERFÜLLEN.

Brewer zeigt mit seinem bekannten **CAP-THEOREM** die **ENTWURFSGRENZEN** für ein VERTEILTES SYSTEM auf. Das **CAP THEOREM** [12][16] dient in diesem Zusammenhang als fundamentaler **FOKUSPUNKT** um die komplementären Größen und deren Beziehungen im Umfeld der Architekturdiskussion VERTEILTER SYSTEME zu verstehen. Gleichzeitig wird eine **KLASSIFIZIERUNG** für VERTEILTE SYSTEME im Datenverarbeitungsbereich gegeben. Die Erkenntnisse des CAP Theorems spielen für die Gestaltung verteilter Datenverarbeitungssysteme damit eine zentrale Rolle. □

KONSISTENZ, VERFÜGBARKEIT UND PARTITIONSTOLERANZ

- **CONSISTENCY** - **KONSISTENZ** aller Netzwerkknoten: Alle Knoten sehen zur selben Zeit dieselben Daten. In verteilten Systemen mit replizierten Daten muß sichergestellt sein, dass nach Abschluß einer Transaktion auch alle Replikate des manipulierten Datensatzes aktualisiert werden.
- **AVAILABILITY** - **VERFÜGBARKEIT**, alle Anfragen an das System können beantwortet werden.
- **PARTITION TOLERANCE** - **PARTITIONSTOLERANZ**, das System arbeitet auch bei Verlust von einzelnen Knoten oder dem Ausfall von Teilsystemen, weiter. Das gesamte System muß stabil laufen unabhängig davon, ob ein Knoten geplant oder ungeplant ausfällt. □

CAP THEOREM - DISKUSSION Es läßt sich leicht erkennen, daß je nach eingesetztem System, entweder die **VERFÜGBARKEIT** oder die **KONSISTENZ** der Daten beim Eintreten eines **PARTITION EVENTS** nicht vollständig gewährleistet werden kann. Es sind damit immer nur 2 der 3 geforderten Eigenschaften **GLEICHZEITIG** erfüllbar.

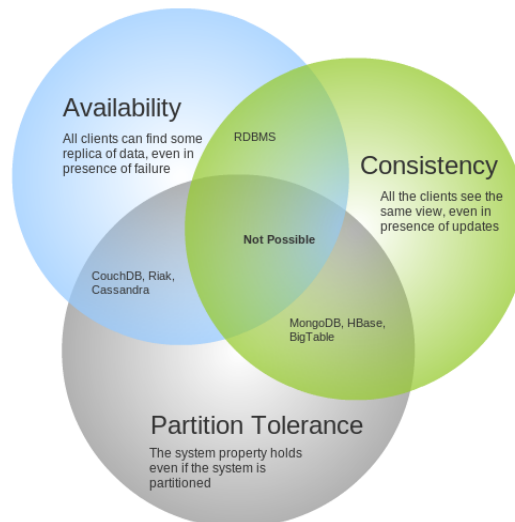
Ein System, das alle drei Eigenschaften zur selben Zeit erfüllt, müsste demnach bei einer Unterbrechung der Verbindung zwischen den Knoten die separierten Knoten verfügbar halten und in diesem Sinne konsistent auf Anfragen reagieren. Wenn nun Daten auf einem der Knoten geändert werden, können offenbar nicht alle restlichen Knoten konsistent gehalten werden, da ja manche Knoten von den anderen Knoten aus nicht

mehr erreichbar sind. Das System kann nun entweder INKONSISTENT oder GAR NICHT REAGIEREN, müsste also entweder **CONSISTENCY** oder **AVAILABILITY** aufgeben. □

KLASSIFIZIERUNG VON APPLIKATIONEN

Da stets nur 2 der 3 Eigenschaften gleichzeitig erfüllt sein können, finden wir eine Klassifizierung von Anwendungen in 3 disjunkte Klassen.

- **CA SYSTEME:** **KONSISTENZ** und **VERFÜGBARKEIT** zur selben Zeit gewährleistet
- **CP SYSTEME:** **KONSISTENZ** und **PARTITIONSTOLERANZ** zur selben Zeit gewährleistet
- **AP SYSTEME:** **VERFÜGBARKEIT** und **PARTITIONSTOLERANZ** zur selben Zeit gewährleistet



CA SYSTEME Beschreibt die Klasse der **RDBMS** Systeme.

- **TRAITS:** **2 PHASE COMMIT, CACHE VALIDATION PROTOCOLS, STRENGE KONSISTENZ**
- **VERZICHT AUF PARTITION TOLERANCE:** Kann **PARTITION TOLERANCE** nicht mehr gewährleistet werden, muß das System auf einer einzelnen Serverinstanz ausgeführt werden.
- **VERTRETER:**
 - **RDBMS SYSTEME**

- SINGEL-SITE DB
- CLUSTER DB
- ...□

CP SYSTEME Beschreibt eine Klasse von **NoSQL** Systemen.

- **VERZICHT AUF AVAILABILITY:** Ereignet sich ein PARTITION EVENT müssen betroffene Services warten, bis Daten konsistent sind und sind erst nachher wieder verfügbar.
- **VERTRETER:**
 - BIGTABLE
 - MONGODB
 - HBASE
 - ...□

AP SYSTEME: Beschreibt eine Klasse von **NoSQL** Systemen.

- **VERZICHT AUF DATENKONSISTENZ**
- **VERTRETER:**
 - CASSANDRA
 - COUCHDB
 - DNS
 - ...□

TECHNOLOGIEENTSCHEIDUNG Bei der Erstellung der PROTOTYPEN stellen wir die **NoSQL SYSTEME** (CP- UND AP SYSTEME) den **RDBMS SYSTEMEN** (CA SYSTEME) gegenüber.

RDBMS vs NoSQL

Wir isolieren das **LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** der PROTOTYPEN als die entscheidenden NICHTFUNKTIONALEN ANFORDERUNGEN in unserer Klasse von Softwareystemen. Ein performantes **LAUFZEITVERHALTEN** in Verbindung mit der Notwendigkeit, **MILLIONEN VON BENUTZERREQUESTS** zur selben Zeit zu bearbeiten, ist die Quadratur des Kreises, die es zu meistern gilt.

	NoSQL Sys- TEME	RDBMS Sys- TEME
SKALIERBARKEIT:	HORIZONTALE SKALIERBARKEIT	VERTIKALE SKALIERBARKEIT
KONSISTENZMODELL:	ACID	BASE
DATENSHEMA:	SCHEMALESS	SCHEMA
DATASTORE:	RDBMS	NoSQL

Figure 1.1: Konzeptvergleich

Eine Vielzahl von Requests gleichzeitig zu bearbeiten, birgt die Notwendigkeit, die Anwendung auf einen CLUSTER von Rechnern zu verteilen. Zur NICHTFUNKTIONALEN Anforderung des LAUFZEITVERHALTENS gesellt sich die Forderung nach einem SKALIERENDEN System.

Das **LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** des Systems sind Größen, die in engem Zusammenhang mit anderen GRÖSSEN und KONZEPTEN stehen. Zur Etablierung unserer Ergebnisse entwickeln wir einen **CP-** und **CA PROTOTYPEN**. Die beiden PROTOTYPEN werden nach grundlegend unterschiedlichen PARADIGMEN und TECHNOLOGIE-KONZEPTEN entwickelt.

SKALIERBARKEIT

SKALIERBARKEIT wird eingesetzt, um die VERFÜGBARKEIT von Anwendungen im industriellen Umfeld gewährleisten zu können. Ein SOFTWARESYSTEM wird als **SKALIERBAR** bezeichnet, wenn das SYSTEM in der Lage ist gleichzeitig zur wachsenden **SYSTEMLAST** den vereinbarten Dienst zu gewährleisten.

Generell wird unterschieden zwischen

- **HORIZONTALE/LINEARE SKALIERUNG:** **HORIZONTALE SKALIERUNG** ist die komplexere der beiden Skalierungsformen. Ein **HORIZONTAL SKALIERENDES** SOFTWARESYSTEM ist in der Lage bei BELIEBIGER SYSTEMLASTEN den DIENST des SYSTEMS zur Verfügung zu stellen. Die Lösung wird nicht durch teurere Hardware erkaufte, sondern durch einen entsprechenden **TECHNOLOGIEENTWURF** erreicht.
BEI DIESER ART DES ENTWURFS WIRD DIE SOFTWARE DES SYSTEMS AN DIE HARDWARE ANGEPASST.
- **VERTIKALER SKALIERUNG:** **VERTIKALE SKALIERUNG** ist die einfachere der beiden Formen der Skalierung. Im Grunde besteht die Lösung darin die Applikation auf leistungsfähigere Computer zu verlagern. Die SKALIERUNG wird damit

eine Frage der Finanzierung. Teurere Hardware erkaufte eine temporäre Lösung mit dem Vorbehalt, daß der erfolgreichen Verarbeitung in letzter Instanz physikalische Grenzen auferlegt sind.

BEI DIESER ART VON SOFTWARESYSTEMEN WIRD DIE HARDWARE DER SOFTWARE ANGEPAßT

KONSISTENZMODELL: ACID vs BASE

Die Entscheidung für ein **KONSISTENZMODELL** ist fundamental für das Verhalten einer Applikation und für die Gestaltung ihrer Architektur.

KLASSIFIZIERUNG VON KONSISTENZMODELLEN

- **STRICT CONSISTENCY:** Alle Leseoperationen liefern den Wert, der die letzte abgeschlossene Schreiboperation geschrieben hat, egal auf welchem Node die Operationen ausgeführt werden.
Bei verteilten Systemen braucht es dafür ein verteiltes Transaktions - Protokoll.
- **EVENTUAL CONSISTENCY:** **EVENTUAL CONSISTENCY** gehört zur Kategorie der schwachen Konsistenzmodelle. Eine Leseoperation liefert eventuell den zuletzt geschriebenen Wert. Je nachdem, an welchem Node der Client die Anfrage schickt, sieht er einen inkonsistenten Zustand.
- **FUNCTIONAL PARTITIONING:** **FUNCTIONAL PARTITIONING** ist eine Möglichkeit eine **skalierbare Lösung** für ein Softwaresystem umzusetzen. **FUNCTIONAL PARTITIONING** besteht darin, das Datenschema als logische Distribution der Daten in funktionale Einheiten zu strukturieren. Daten die in funktionaler Weise zusammenhängen, werden dabei in einer Serverinstanz gruppiert. Die Datenverarbeitung innerhalb von Transaktionen kann in diesem Sinne schnell und effizient abgewickelt werden. □

ACID Je nachdem für welche Art von **KONSISTENZ** wir uns entscheiden können wir eine Zahl von Eigenschaften fordern. Entscheiden wir uns für **STRICT CONSISTENCY**, werden die **ACID** Eigenschaften gefordert:

- **ATOMICITY:** ABGESCHLOSSENHEIT - Jede Operation innerhalb einer Transaktion wird vollständig abgeschlossen (oder alle Operationen werden wieder rückgängig gemacht)
- **CONSISTENCY:** KONSISTENZ - Bei Transaktionsstart und -ende sind die Daten in einem **KONSISTENTEN** Zustand.

- **ISOLATION:** ABGRENZUNG - Die TRANSAKTION verhält sich so, als wäre sie die einzige Operation, die auf dem Datenbestand ausgeführt wird.
- **DURABILITY:** DAUERHAFTIGKEIT - Ist eine TRANSAKTION einmal abgeschlossen, sind die Veränderungen dauerhaft im System gespeichert.

Ein KONSISTENZMODELL, das die **ACID** Eigenschaften erfüllt, ist in der Lage, komplexe Datenschemen und deren Abhängigkeiten abzubilden. □

BASE **WEB 2.0** Systeme müssen in erster Linie performant und zuverlässig sein.

Eigenschaften:

- **BASICALLY AVAILABLE:** Die Applikation ist grundsätzlich immer verfügbar
- **SOFT STATE:** Die Applikation muss nicht immer in einem **KONSISTENTEN ZUSTAND** sein
- **EVENTUAL CONSISTENCY:** Nach dem Verständnis von **BASE** wird die **KONSISTENZ** der Daten als ein Zustand betrachtet, der irgendwann erreicht wird. Das ist die Idee des **EVENTUAL CONSISTENCY**. Es wird damit in Kauf genommen, dass sich bis zum Erreichen der Konsistenz, die Daten in einem **INKONSISTENTEN ZUSTAND** befinden.

Die HOCHVERFÜGBARKEIT der Daten geht zu Lasten der **KONSISTENZ**. Die Daten sind im Anschluß an einen Schreibvorgang nur **EVENTUELL KONSISTENT**. Erst nach dem Verstreichen einer gewissen Zeit wird die Änderungen an alle anderen REPLIKATE weitergegeben. □

VERGLEICH ACID vs BASE

ACID	BASE
STRENGE KONSISTENZ	SCHWACHE KONSISTENZ
ISOLATION	VERFÜGBARKEIT
FOKUS AUF COMMIT	INETWA RICHTIGE ANTWORTEN
VERSCHACHTELTE TRANSAKTIONEN	SCHNELL
KONSERVATIV/PESSIMISTISCH	AGGRESSIV/OPTIMISTISCH
KOMPLIZIERTE EVOLUTION	EINFACHE EVOLUTION

DATENVERARBEITUNG

DB VERTEILTE DATENBANKSYSTEME kommen zum Einsatz wenn ein Rechner allein, als Datenbankserver, die anfallenden Aufgaben nicht mehr bewältigen kann.

STAND ALONE DATENBANKSYSTEME stoßen schnell an ihre Grenzen falls:

- **VOLUMEN:** Das VOLUMEN der zu speichernden Daten eine kritische Grenze übersteigt
- **ZAHL:** Die Zahl der Requests an den DB SERVER eine kritische Grenze übersteigt
- **REPLIKATION:** Die Daten im Sinne von DATENSICHERHEIT über mehrere Rechner repliziert werden □

DATENVERARBEITUNG Vertreter für die, für uns interessante Klasse von SOFTWARESYSTEMEN müssen in der Lage sein, in kürzester Zeit immense Datenmengen **ZU BEWEGEN** und **ZU VERARBEITEN**. Wir erkennen, daß die **DATENVERARBEITUNG** eine kritische Größe für das Verständnis der Anforderungen an den **TECHNOLOGIEENTWURF** darstellt. Finden wir eine effiziente Architektur, können wir **SALIERBARKEIT** und ein **PERFORMANTES LAUFZEITVERHALTEN** garantieren. □

NoSQL Eine Entwicklung im DATENVERARBEITUNGSBEREICH, bekanntgeworden als **NoSQL** Datenbanken, wird vorangetrieben, um ein performantes **LAUFZEITVERHALTEN** und **SKALIERBARKEIT** im Industriebereich zu schaffen. Die grundlegende Idee ist hier, die Restriktionen RELATIONALER DATENBANKEN zu umgehen und andere Wege in der Datenverarbeitung zu beschreiten.

Die Abkehr von etablierten DATENVERARBEITUNGSTECHNIKEN hin zu neuen, unkonventionellen Lösungen steht im Mittelpunkt dieses Ansatzes. **ACID** als Ankerpunkt des **KONSISTENZMODELL** wird durch die neu formulierten **BASE** Kriterien abgelöst. **NoSQL** verfolgt den Ansatz daß eine zeitliche Lockerung der KONSISTENZKRITERIEN die gewünschten Ergebnisse bringt, ohne geforderte Softwareeigenschaften zu verletzen. □

FRAGESTELLUNG DER ARBEIT

VERTIKALE SKALIERUNG VS. HORIZONTALE SKALIERUNG Die bestimmende technologische Herausforderung der letzten Jahre in der SOFTWAREENTWICKLUNG war die Eskalierung der zu bearbeitenden Datenmengen in einer entsprechenden vorgegebenen Zeit. Etablierte VERTEILTE RDBMS Systeme stossen in diesem Zusammenhang zunehmend an ihre Grenzen. Neue Konzepte müssen entwickelt werden, um den geänderten Anforderungen gerecht zu werden.

Die grundlegende Frage in diesem Sinne ist, ob **VERTIKAL SKALIERENDE** Systeme in diesem Bereich, verworfen werden müssen und durch **HORIZONTAL SKALIERENDE** Lösungen ersetzt werden sollen? □

FRAGESTELLUNG Die grundlegende Frage in diesem Zusammenhang ist damit:

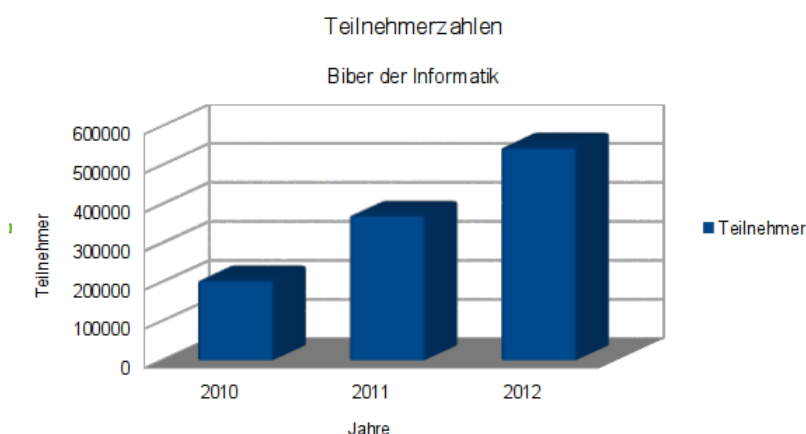
SOLL ZUR ETABLIERUNG EINES VERTEILTEN WEBBASIERTEN CLIENT-SERVER SYSTEMS MIT MILLIONEN VON GLEICHZEITIGEN ANWENDERZUGRIFFEN UND DER NOTWENDIGKEIT DATEN IM PETA BYTE BEREICH ZU BEARBEITEN BEI GLEICHZEITIG PERFORMANTEM LAUFZEITVERHALTEN, EIN VERTIKAL ODER EIN HORIZONTAL SKALIERENDES SYSTEM EINGESETZT WERDEN? □

DAS BIBER SOFTWARESYSTEM - EINE BESCHREIBUNG

BESCHREIBUNG UND GESCHICHTE DES “BIBERS“

Der **INFORMATIK BIBER** (**BEBRAS INTERNATIONAL CONTEST ON INFORMATICS AND COMPUTER LITERACY**) ist ein jährlich stattfindender Informatik Wettbewerb. Der Wettbewerb wird durchgeführt in Volks- und Mittelschulen. Am Computers lösen die Teilnehmer 15 bis 21 Aufgaben. Die zu lösenden Aufgaben umfassen dabei interaktive und Multiplechoice Aufgaben.

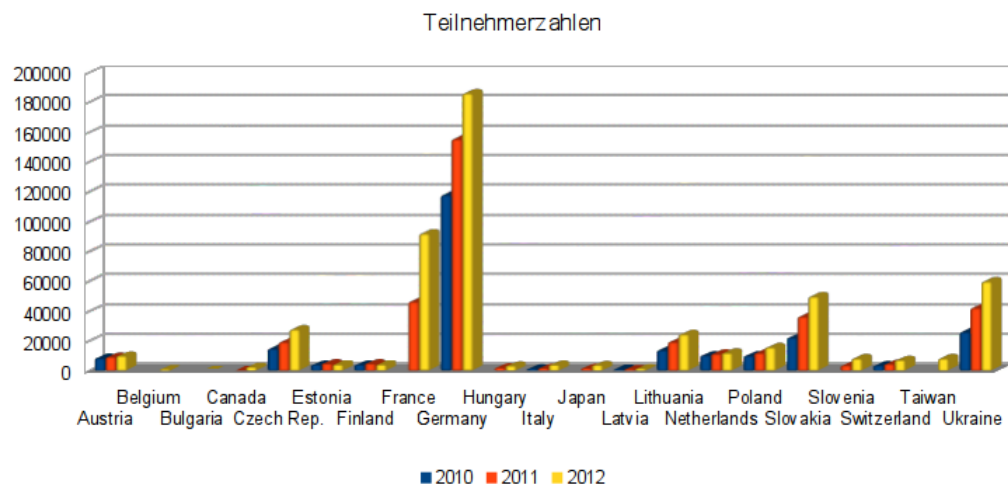
Die Zahl der am Wettbewerb teilnehmenden Länder wächst jährlich. Mit der zunehmenden Zahl an Teilnehmern wächst die Notwendigkeit nach einer skalieren Lösung für das System.



GESCHICHTE DES BIBERS

- Der Wettbewerb wurde zum ersten Mal **2004** unter dem Namen Bebras in Litauen durchgeführt.
- Im Jahre **2005** wurde Beaver in
 - Estland,
 - Lettland,
 - Litauen,
 - den Niederlanden,

- Polen durchgeführt.
- **2006** wurde der Informatik-Biber im Rahmen des Informatikjahres unter dem Namen EI:SPIEL blitz! erstmals in Deutschland durchgeführt.
- **2007** fand der Informatik-Biber in Deutschland, Österreich, Polen, Litauen und den Niederlanden statt.
- Im Jahr **2010** fand der Informatik-Biber zum ersten Mal auch in der Schweiz statt.



USECASEBESCHREIBUNG

Zur transparenten Darstellung der Abhängigkeiten und Eigenschaften des Biber Softwaresystems sollen die Usecases des Systems näher erörtert werden. Das Verständnis der Aktoren, deren Rechte und Abhängigkeiten, sollen uns helfen die Anforderungen an das System besser zu verstehen.

AKTOREN-RECHTE Im folgenden finden wir das Zusammenspiel der Rechte und Aktoren

- siehe Abbildung 1.2
- siehe Abbildung 1.3
- siehe Abbildung 1.4

ID	ACRONYM	BESCREIBUNG
1.1	CPA	Anlegen eines Proxy-Administrators
1.2	UPA	Angleichen eines Proxy-Administrators
1.3	DPA	Löschen eines Proxy-Administrators aus dem System
2.1	CTA	Anlegen eines Test-Administrators
2.2	UTA	Angleichen eines Test-Administrators
2.3	DTA	Löschen eines Test-Administrators aus dem System
3.1	CSA	Anlegen eines School-Administrators
3.2	USA	Angleichen eines School-Administrators
3.3	DSA	Löschen eines School-Administrators aus dem System
4.1	CS	Anlegen einer Schule
4.2	US	Angleichen der Daten einer Schule
4.3	DS	Löschen einer Schule aus dem System
5.1	CP	Anlegen eines Schülers
5.2	UP	Angleichen der Daten einer Schülers
5.3	DP	Löschen eines Schülers aus dem System
6.1	CSS	Anlegen von Schulverwaltungsdaten
6.2	USS	Angleichen von Schulverwaltungsdaten
6.3	DSS	Entfernen von Schulverwaltungsdaten aus dem System
7.1	CSS	Anlegen von Schulverwaltungsdaten
7.2	USS	Angleichen von Schulverwaltungsdaten
7.3	DSS	Entfernen von Schulverwaltungsdaten aus dem System
8.1	CE	Anlegen einer Aufgabe
8.2	UE	Angleichen einer Aufgabe
8.3	DE	Entfernen einer Aufgabe aus dem System
9.1	CC	Anlegen einer Testkategorie
9.2	UC	Angleichen einer Testkategorie
9.3	DC	Entfernen einer Testkategorie aus dem System
10.1	CE	Anlegen einer Testevaluierung
10.2	UE	Angleichen einer Testevaluierung
10.3	DE	Entfernen einer Testevaluierung aus dem System
11.1	CT	Anlegen einer Tests
11.2	UT	Angleichen eines Tests
11.3	DT	Entfernen eines Tests aus dem System
12.1	ST	Test ausführen

Figure 1.2: Aktorenrechte

AKTOR	RECHTE
System-Administrator	{CPA, UPA, DPA}, {CTA, UTA, DTA}, {CSA, USA, DSA}, {CS, US, DS}, {CP, UP, DP}, {CSS, USS, DSS}, {CE, UE, DE}, {CC, UC, DC}, {CE, UE, DE}, {CT, UT, DT}, ST
Proxy-Administrator	{CTA, UTA, DTA}, {CSA, USA, DSA}, {CS, US, DS}, {CP, UP, DP}, {CSS, USS, DSS}, {CE, UE, DE}, {CC, UC, DC}, {CE, UE, DE}, {CT, UT, DT}, ST
Test-Administrator	{CSA, USA, DSA}, {CS, US, DS}, {CP, UP, DP}, {CSS, USS, DSS}, {CE, UE, DE}, {CC, UC, DC}, {CE, UE, DE}, {CT, UT, DT}, ST
School-Administrator	{CS, US, DS}, {CP, UP, DP}, {CSS, USS, DSS}, {CE, UE, DE}, {CC, UC, DC}, {CE, UE, DE}, {CT, UT, DT}, ST
Test Participant	ST

Figure 1.3: Aktorenliste

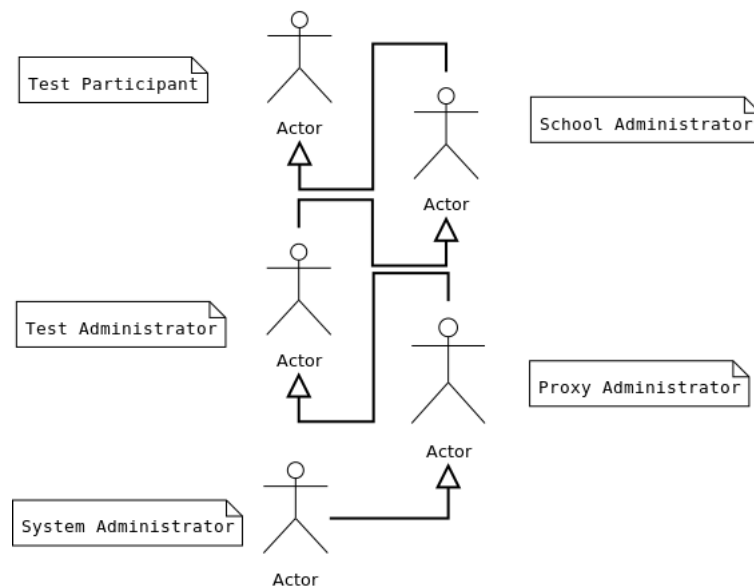


Figure 1.4: Beziehung unter Aktoren

SYSTEMBESCHREIBUNG - ANFORDERUNGEN

Die für die Durchführung des Wettbewerbs notwendigen technischen Restriktionen sollen kurz besprochen werden.

WEBTECHNOLOGIE Das System hat den vorgegebenen technischen infrastrukturellen Unzulänglichkeiten der pädagogischen Einrichtungen Rechnung zu tragen. Die Herausforderung besteht darin die Applikation auf einer Zahl von rudimentären Internet-Technologien aufzubauen.

□ **HTTP: Hyper Text Transfer Protocol** (Quelle:[18])

Ist ein Protokoll zur Übertragung von Daten über ein Netzwerk. Es wird hauptsächlich eingesetzt, um Webseiten aus dem World Wide Web in einen Webbrowser zu laden.

HTTP gehört der sogenannten Anwendungsschicht etablierter Netzwerkmodelle an. Die Anwendungsschicht wird von den Anwendungsprogrammen angesprochen, im Fall von HTTP ist das meist ein Webbrowser. Im ISO/OSI-Schichtenmodell entspricht die Anwendungsschicht den Schichten 5–7. HTTP ist ein zustandsloses Protokoll.

Ein zuverlässiges Mitführen von Sitzungsdaten kann erst auf der Anwendungsschicht durch eine Sitzung über eine Session-ID implementiert werden.

□ **CLIENT TECHNOLOGIE:** Client-Technologie sind lediglich bezüglich des Einsatzes von Plugins beschränkt.

Gängige Browser werden durchgehend unterstützt.

- **FIREFOX**
- **CHROME**
- **OPERA**
- **INTERNET EXPLORER**
- ...

□ **HTML 5** (Quelle:[19])

HTML 5 ist eine textbasierte Auszeichnungssprache zur Strukturierung und semantischen Auszeichnung von Inhalten wie Texten, Bildern und Hyperlinks in Dokumenten.

HTML 5 ist als Sammlung mehrerer Technologien, die im Tandem arbeiten, zu verstehen.

- **HTML**
- **JAVASCRIPT**

– CSS

□ **SERVERTECHNOLOGIE**

Die Wahl der einzusetzenden Severtechnologie obliegt keinen Restriktionen.

Die geforderten Restriktionen sollen gewährleisten, daß auch Schulen mit veralteter Ausstattung am Wettbewerb teilnehmen können.

Anforderungsanalyse

KLASSE VON SOFTWAREANWENDUNGEN

Die für uns interessante **KLASSE VON SOFTWAREANWENDUNGEN** möchten wird durch die für sie **CHARAKTERISTISCHEN EIGENSCHAFTEN** beschreiben. Die gewünschte Klasse findet sich in der **DURCHSCHNITTSMENGE** der durch folgende **EIGENSCHAFTEN** und **ANFORDERUNGEN** beschriebenen Softwaresysteme:

- **VERTEILTES SYSTEM**
- **BASIEREND AUF DEM HTTP PROTOKOLL**
- **UNTERSTÜTZT HTML 5**
- **HOCH SKALIERBAR**
- **DATENVERARBEITUNG IM TERRABYTE BEREICH**
- **PERFORMANTES LAUFZEITVERHALTEN** □

VERTEILTE SYSTEME

Die BIBER WEBAPPLIKATION ist eine verteilte Softwareanwendung. Sprechen wir von einem **VERTEILTEN SYSTEM** das Daten verarbeitet, müssen wir einige Eigenschaften fordern: (siehe [10])

1. **HOHE VERFÜGBARKEIT**
2. **ORTSTRANSparenZ**
3. **REPLIKATIONSTRANSparenZ**
4. **VERTEILTE ANFRAGEBEARBEITUNG**
5. **VERTEILTE TRANSPARENZBEARBEITUNG**
6. **HARDWAREUNABHÄNGIGKEIT**
7. **BETRIEBSSYSTEMUNABHÄNGIGKEIT**
8. **LEISTUNGSTRANSparenZ** □

WORLD WIDE WEB

HTTP PROTOKOLL - CLIENT/SERVER Haben wir es mit einem System zu tun, das im Bereich des Internets zum Einsatz kommen soll, müssen zusätzlich zu den Vorgaben eines verteilten Systems weitere Eigenschaften gefordert werden.

Sehen wir uns mit der Herausforderung konfrontiert, ein CLIENT-SERVER System für das Internet zu entwickeln, gibt es eine Reihe von technischen Einschränkungen, die wir erfüllen müssen. In erster Linie leiten sich diese Einschränkungen aus den Eigenschaften des **HTTP PROTOKOLLS** ab. □

EIGENSCHAFTEN DES HTTP PROTOKOLLS

- **ZUSTANDSLOS**
- **REQUEST-RESPONSE KOMMUNIKATION**
- **CLIENT-SERVER ARCHITEKTUR**
- **MIME DATENTYPEN VERSTÄNDNIS** □

ANWENDUNGSANFORDERUNGEN

WEBANWENDUNG Die BIBER WEBAPPLIKATION selbst stellt einige grundlegende Anforderungen an die Gestaltung der Prototypen. Der gleichzeitige Zugriff mehrerer tausend Teilnehmer auf den Dienst der Webapplikation resultiert in der Formulierung kritischer Anforderungen an das System.

- **WEBANWENDUNG:** Die Anwendung ist als CLIENT-SERVER System für das Internet umzusetzen. In diesem Zusammenhang müssen wir alle Restriktionen beachten die bei der Gestaltung des Prototyps im Umfeld des Internets zu befolgen sind.
- **SKALIERBARKEIT:** Der Spezifikation entsprechend muß das System in der Lage sein, tausende von Anfragen gleichzeitig zu verarbeiten. Von Jahr zu Jahr steigt die Zahl der am Wettbewerb teilnehmenden Schulen. Gleichzeitig finden immer mehr Länder den Weg in den Wettbewerb.

Die Prototypen müssen in der Lage sein, zu skalieren, um die anfallende SERVER-LAST zu bewältigen.

- **PERFORMANTE DATENVERARBEITUNG:** Die PROTOTYPEN müssen in der Lage sein, in kürzester Zeit immense Datenmengen zu bewegen. Wir erkennen, daß die DATENVERARBEITUNG eine kritische Größe für das Verständnis der Anforderungen an die Architektur darstellt. Finden wir eine effiziente Architektur, können wir **SKALIERBARKEIT** und ein performantes LAUFZEITVERHALTEN garantieren.

Eine ausführliche Diskussion zu diesem Thema soll uns helfen, ein Verständnis für die geeignete Lösung zu etablieren.

- **REPLIKATIONSFÄHIGKEIT:** Aus der Forderung nach **SKALIERBARKEIT** des Systems erwächst die direkte Forderung nach einem Replikationsmechanismus innerhalb des Systems. REPLIKATION ist eine der Technologien, um SKALIERBARKEIT zu etablieren.
- **PERFORMANTES LAUFZEITVERHALTEN:** Das **LAUFZEITVERHALTEN** einer Applikation ist das direkte Resultat der gewählten **REPLIKATIONSSTRATEGIE**, des **KONSISTENZMODELLS** und des **DATENSCHEMAS**. □

NoSQL: Etablierung eines neuen Datenverarbeitungs Standards

EINFÜHRUNG

DEFINITION: NoSQL Bisher gibt es keine **EINHEITLICHE** und **ABGRENZENDE** Definition für **NoSQL DATENBANKENSYSTEME**, dafür ist das Technologieparadigma noch zu neu.

Hier jedoch ein Versuch einer Definition wie sie von Edlich [7] vorgenommen wurde:

- **KEIN RELATIONALES DATENMODELL:** Postuliert wird eine Vielzahl neuer Modelle mit dem Gedanken, daß das **RELATIONALE MODELL** in bestimmten Bereichen an seine Grenzen stößt oder nicht zweckmäßig ist.
- **EIGNUNG FÜR SYSTEME MIT VERTIKALER UND HORIZONTALER SKALIERBARKEIT:** Schnelle Reaktionszeiten auf **ANFRAGEN** und **MANIPULATIONEN** lassen sich nur mit dem **SCALE-OUT** Prinzip erreichen, also mit **HORIZONTALER SKALIERUNG** durch DYNAMISCHES EINBINDEN/LÖSCHEN von Rechnerknoten (*Nodes*). Die gängige Praxis für RDBS hingegen ist das **SCALE UP** Prinzip, bei dem ein Rechner immer weiter technisch aufgerüstet wird.
- **SCHEMAFREI:** Zur Bewältigung der Datenverwaltungsaufgaben mit ständig wachsendem Datenvolumen, in kürzest möglicher Zeit wird die SCHEMAINFORMATION in die Anwendung verlagert weg vom Datenbanksystem.

- **EINFACHE DATENREPLIKATION ZUR UNTERSTÜTZUNG DER VERTEILTEN ARCHITEKTUR:** Die Grundidee der Verteilung des Datenbestands auf viele Knoten erfordert eine einfache und schnelle **REPLIKATION** der Daten. Die Daten müssen auch bei Knotenausfällen weiter verfügbar sind. Möglich wird dies durch „abgeschwächte“ Konsistenzanforderung. Wir sehen eine praktische Umsetzung im **BASE-KONSISTENZMODELL**.
- **EIN ANGEPASTES KONSISTENZMODELL:** Kein **ACID** als Konsistenzmodell
□

TECHNOLOGIEPARADIGMEN

NoSQL Datenbanken bieten eine Fülle von Eigenschaften, die für die Entwicklung VERTEILTER, DATENVERARBEITENDER Systeme vorteilhaft sind. Einige der Eigenschaften sind:

- **GUTES LAUFZEITVERHALTEN**
- **SKALIERBARKEIT**
- **EFFIZIENTES REPLIKATIONSMODELL und PARTITION TOLERANCE**

Diese Eigenschaften werden erreicht durch die ETABLIERUNG unterschiedlicher **TECHNOLOGIEPARADIGMEN**. Die wichtigsten dieser **TECHNOLOGIEPARADIGMEN** sollen hier näher beschrieben werden. □

CONSISTENT HASHING

CONSISTENT HASHING FUNKTIONEN **CONSISTENT HASHING FUNKTIONEN** als zentrales Konzept von **NoSQL** Systemen verstehen sich als eine Weiterentwicklung von HASHING FUNKTIONEN im Datenverarbeitungsbereich.

Das Konzept der **CONSISTENT HASHING FUNKTIONEN** trägt dem Umstand Rechnung daß eine Datenbank über mehrere tausend Rechner verteilt sein kann. Die Herausforderung besteht in diesem Zusammenhang darin, **KNOTENZUGÄNGE** wie auch **KNOTENNAUSEFÄLLE** zu kompensieren, ohne daß der laufende Betrieb beeinträchtigt wird. □

FUNKTIONSWEISE Implementierung von **CONSISTENT HASHING FUNKTIONEN** siehe [14]

- Der Wertebereich der Zielmenge wird nicht als Liste, sondern als Ring verstanden. Es werden beide Enden zusammengefügt
- Die **SERVER** 1 - n werden gemäß des HASHWERTES ihrer **ID's** (**SERVERNAMEN**, **IP-ADRESSEN**, ...) auf diesem Adressring angeordnet.
- Die Datensätze werden gemäß ihres HASHWERTES auf dem Adressring angeordnet. Die Datensätze werden auf den **SERVERN** gespeichert, die im Uhrzeigersinn auf dem Adressring nachfolgend angeordnet sind.
- Wird ein neuer **SERVER** hinzugefügt, so wird er gemäß seines HASHWERTES auf dem Ring angeordnet.

Die Datensätze, deren Adressen nun im Uhrzeigersinn vor dem neuen Server liegen, werden vom ursprünglichen Server auf den neuen Server umgespeichert. Alle anderen Datensätze können an der ihnen zugewiesenen Adresse bleiben.

- Wird ein **SERVER** herausgenommen, dann werden seine Daten auf den ihm Uhrzeigersinn nachfolgenden Server kopiert. Alle anderen Datensätze sind nicht betroffen von dieser Änderung.
- Verfügt ein Server über mehr KAPAZITÄT, so können für ihn entsprechend viele **VIRTUELLE SERVER** erzeugt werden.

Diese **VIRTUELLEN SERVER** unterscheiden sich in der ID und damit werden für jeden dieser **VIRTUELLEN SERVER** andere HASHWERTE ermittelt und damit auch unterschiedlich auf dem Adressring plaziert.

Auf den realen Server werden dann alle Daten kopiert, deren Adressen von den zugehörigen **VIRTUELLEN SERVERN** auf dem Ring angeordnet sind. □

EINSATZ **CONSISTENT HASHING** dient als wichtige Grundlage für im **NoSQL** Bereich eingesetzten **DATENMODELLE**. **CONSISTENT HASHING** wird eingesetzt um **PARTITION TOLERANCE** und **LINEARE SKALIERBARKEIT** implementieren zu können. □

MVCC - MULTIVERSION CONCURRENCY CONTROL

Das **MULTIVERSION CONCURRENCY CONTROL VERFAHREN** stellt sicher, daß konkurrierende Zugriffe auf Datensätze durch die Verwaltung verschiedener, unveränderlicher Versionen der Datensätze kontrolliert werden.

Für jeden datenändernden Zugriff

- **EINFÜGEN**
- **ÄNDERN**
- **LÖSCHEN**

wird für den Datensatz eine neue Version erstellt. Die verschiedenen Versionen werden in einer zeitlichen Reihenfolge gereiht als Liste von Versionen. □

FUNKTIONSWEISE **MVCC** stellt ein wichtiges Konzept für **NoSQL** Systemen dar, das es ermöglicht, konkurrierende Zugriffe auch ohne das **SPERREN** von Datensätzen zu koordinieren.

- **LESEN:** Das Lesen ist vom Schreiben völlig entkoppelt. Zu jedem Zeitpunkt steht eine aktuelle Version zum Lesen bereit.
- **SCHREIBEN:** Jeder Prozess, der die aktuelle Version ändert, erstellt eine neue Version dieses Datensatzes mit einem Verweis auf die gelesene aktuelle Version.

Beim Transaktionsende wird die Vorgänger-Versionsnummer des aktuell in dieser Transaktion geänderten Datensatzes mit seiner aktuellen Versionsnummer verglichen.

- **AUFRÄUMEN:** In periodischen Zeitabständen müssen alte Versionen gelöscht werden. Ein Datensatz wird als alt eingestuft, wenn er von keiner Transaktion mehr verwendet wird. □

EINSATZ MVCC als **TECHNOLOGIE** kommt in den meisten **NOSQL DATENMODELLEN** zum Einsatz. Gleichzeitig greifen auch einige **RDBMS** Systeme auf diese **TECHNOLOGIE** zurück. □

MAP/REDUCE

MAP/REDUCE ist weniger ein **TECHNOLOGIEPARADIGMA** als ein **ALGORITHMUS** zur **VERTEILTEN DATENVERARBEITUNG**. MAP/REDUCE wird in 2 Hauptphasen durchgeführt:

- **MAPPING PHASE**
- **REDUCTION PHASE**

ALGORITHMUS MAP/REDUCE ist ein **VERTEILTER ALGORITHMUS**, der darauf ausgelegt ist, mit der Zahl der verfügbaren **RESSOURCEN** linear zu skalieren. Voraussetzung für einen möglichen Einsatz des Verfahrens ist daß Abhängigkeiten innerhalb der Eingabedaten nicht hinderlich sind für eine Aufteilung der Daten in Blöcke beliebiger Größe. Die Zahl der Blöcke skaliert in bijektiver Weise mit der verfügbaren Ressourcen.

In jeder der Phasen werden die Daten einem bestimmten Ablauf folgend verarbeitet.

- **PARTITIONING:** Die Eingabedaten in Form einer List als **KEY/VALUE** Paare vorliegen. Die Daten werden nun sinnvoll über die zur Verfügung stehenden Rechner verteilt.
- **MAPPING PHASE:** In der **MAPPING PHASE** leitet der **MAP/REDUCE** Algorithmus die Eingabedaten an den **MAPPER** weiter.
Die Aufgabe des Mappers ist es die Daten zu filtern und in einen sinnvolles Format zu transformieren. Der Vorgang der Transformation muß in sich abgeschlossen sein und darf für jedes **KEY/VALUE** Paar nur vom Ausgangszustand des gegebenen Paares abhängig sein.
Die Ausgabedaten der **MAP PHASE** dienen als Eingabe für die 2te Phase des Algorithmus. Die Daten liegen weiter in Form einer List als **KEY/VALUE** Paare vor.
- **SHUFFLING:** Während des **SHUFFLINGS** der transformierten Daten, werden jene Daten, die denselben **KEY** der Vorgabe der Applikation entsprechend zusammengefügt und anschließend an den **REDUCER** weitergeleitet.

- **REDUCE PHASE:** Die Aufgabe des REDUCERS ist es, auf der gesamten Datenmenge die finalen Schritte auszuführen. □

QUELLEN

- siehe [7]
- siehe [13]
- siehe [14]

NoSQL DATENMODELLE

Konkrete Implementierungen von **NoSQL DATENBANKEN** vereinen in der Regel mehrere mögliche DATENMODELLE. Die DATENMODELLE und die ihnen zugrundeliegenden Eigenschaften sind darauf abgestimmt, als Einheit diverser Technologien zu arbeiten.

KEY/VALUE DATENMODELL

KEY/VALUE DATENMODELLE bedienen sich eines einfachen **KEY/VALUE** Systems. Ein KEY ist immer einem bestimmten SCHLÜSSEL zugeordnet, der aus einer **STRUKTURIERTEN** oder **WILLKÜRLICHEN** Zeichenkette bestehen kann.

EIGENSCHAFTEN Das KEY-VALUE DATENMODELL hat eine Reihe erfreulicher Eigenschaften, die es von anderen Datenmodellen abhebt:

- **LINEAR SKALIERBAR**
- **PARTITIONIERBAR**
- **HOCH VERFÜGBAR**
- **KURZE ANTWORTZEITEN**

Die hohe **VERFÜGBARKEIT** der Daten wird durch ein **KONSISTENZMODELL** erkauft das nur die **BASE** Kriterien erfüllt. Einschräkend kommt hinzu daß nur einfache Datenstrukturen in einem KEY/VALUE DATENMODELL abgebildet werden können. □

TECHNOLOGIEN

- **CONSISTENT HASHING**
- **MVCC** □

SPALTENORIENTIERTES DATENMODELL

Der Ansatz, Daten nicht in **ZEILEN-** (**RELATIONALE DATENBANKEN**) sondern in **SPALTENFORM** zu speichern, ist ein fundamentales Konzept im **NoSQL** Bereich.

STRUKTUR Ein SPALTENORIENTIERTES Design der Datenstruktur ermöglicht, daß zusammengehörende Daten in einem **DISK BLOCK** abgelegt werden können. Ein Eintrag besteht dabei aus:

- **NAME DER SPALTE**
- **DATEN**
- **ZEITSTEMPEL.** □

COLUMN FAMILY Spalten mit **ÄHNLICHEN** oder **VERWANDTEN** Inhalten bilden dabei die sogenannte COLUMN FAMILY. In ihrem Aufbau entspricht die COLUMN FAMILY einer Tabelle im **RELATIONALEN MODELL**. Im Gegensatz zur Tabelle hat die COLUMN FAMILY keine strikte Struktur. Sie kann aus Tausenden oder sogar Millionen von Spalten bestehen. □

ZELLE Die kleinste Datenverarbeitungseinheit ist eine ZELLE. Spaltenorientierte Datenmodelle mit VERSIONING sind in Lage, mehrere Versionen desselben Datensatzes in einer Zelle zu verwalten. □

EIGENSCHAFTEN Durch die Speicherung der Daten in Spaltenform erkaufen wir uns einige willkommene Eigenschaften, die im **RELATIONALEN** Modell nicht gegeben sind:

- Im Zuge des Leseprozesses wird nur die gewünschte Information gelesen
- Der Schreibprozess besteht in der Regel aus einem einzelnen **DISK SEEK**

Einzig im Fall daß eine Vielzahl von Daten zu schreiben ist, die über mehrere Spalten verstreut sind, bemerken wir eine Verlangsamung im Schreibprozess. □

DOKUMENTORIENTIERTES DATENMODELL

Das DOKUMENTORIENTIERTE DATENMODELL stellt ein fundamentales Konzept im Bereich der **NoSQL** Datenbanken dar.

DOKUMENT Im Bereich der DOKUMENTORIENTIERTEN DATENBANKEN werden Daten als Dokumente abgelegt. Ein DOKUMENT versteht sich in diesem Sinne als eine strukturierte Zusammenstellung der Daten. Der Begriff des Dokuments sollte in diesem Zusammenhang nicht verwechselt werden mit den von Textverarbeitungsprogrammen erzeugten Dateien. □

DATENSTRUKTUR Die Daten in einem Dokument sind in Form von KEY/VALUE-PAAREN strukturiert. Das DOKUMENTORIENTIERTE DATENMODELL genügt damit keiner Normalform, eine spezifische Strukturierung der zu speichernden Daten ist nicht vorgeschrieben. In jedem einzelnen Dokument kann eine andere Struktur zu finden sein, ein allumfassendes, die ganze Datenbank umspannendes **DATENBANKSCHEMA** existiert nicht.

Das Datenmodell wird als **SCHEMA FREI** bezeichnet. Die Verwendung von **CONSTRAINTS** oder **TRIGGERN** ist im Datenbanksystem nicht möglich. Jedes Dokument stellt für sich gesehen eine geschlossene Einheit mit eigenem Schema dar. Das Datenmodell selbst unterstützt von Haus aus keine Abfragesprache im Sinne von **SQL**. □

EIGENSCHAFTEN Die Stärke DOKUMENTENBASIERTER DATENBANKEN ist es, zusammengehörige Daten als Einheit zu speichern. Im Gegensatz zu **RELATIONALEN DBMS**, die mit **SQL** über einen einheitlichen Standard verfügen, unterscheiden sich die Implementierungen dokumentenbasierter Datenbanken mitunter erheblich in ihrem Funktionsumfang. □

NoSQL IMPLEMENTIERUNGEN

Wir wollen einige Vertreter von **NoSQL DATENBANKSYSTEMEN** näher analysieren, um ihren Wert für einen Einsatz im **NoSQL PROTOTYPEN** abzuschätzen:

- **CASSANDRA**
- **COUCHDB**
- **HADOOP-HIVE**
- **MONGODB**
- **APPEENGINE**

CASSANDRA

BESCHREIBUNG CASSANDRA wurde mit dem Augenmerk entwickelt, große Datenmengen, die über mehrere Netzwerkknoten verteilt sind, zu verwalten.

Die Datenbankarchitektur wurde von vornherein darauf ausgelegt, Knotenausfälle zu kompensieren. Die Daten sind als **REPLIKATE** auf mehreren Knoten vorhanden. Damit wird **RELIABILITY** und **FAULT TOLERANCE** für das System etabliert. □

CASSANDRA CLUSTER Eine CASSANDRA INSTANZ ist eine Sammlung unabhängiger Knoten, die in Clustern zu einer Einheit gebündelt werden.

Alle **KNOTEN** in einem CASSANDRA CLUSTER sind **PEERS**, kein Knoten besitzt mehr **RECHTE** oder **AUFGABEN**. In einem CASSANDRA CLUSTER findet sich damit kein **SINGLE POINT OF FAILURE**. □

EIGENSCHAFTEN

- Die Daten werden über alle Knoten im Netz verteilt.
- Knoten kommunizieren über ein **PEER TO PEER GOSSIP** Protokoll untereinander.
- CASSANDRA basiert auf einem **SPALTENORIENTIERTEN DATENMODELL**. Als Adressierungsmethode wird **CONSISTENT HASHING** eingesetzt.
- Alle Schreibzugriffe auf die Daten werden automatisch im gesamten Cluster repliziert. □

CQL Als Datenmanipulationssprache wird **CQL** eingesetzt. Zur Vereinfachung der Arbeit mit **CASSANDRA** verwendet **CQL** eine ähnliche Syntax wie **SQL**.

Für jede Anfrage kann in **CASSANDRA** der Grad der gewünschten **KONSISTENZ** angegeben werden. Eine Anfrage, die mehr Datenkonsistenz verlangt, ist auch langsamer in der Ausführung. □

KONSISTENZMODELL:

□ **STUFEN DER SCHREIBKONSISTENZ:**

- **ANY:** Gewährleistet, daß die Daten in mindestens einem Knoten gespeichert sind.
- **ONE:** Gewährleistet, daß die Daten in den **COMMIT-LOG** von mindestens einem **REPLICA** gespeichert sind.
- **QUORUM:** Gewährleistet, daß die Daten in einem **QUORUM** von **REPLICAS** gespeichert sind.
- **ALL:** Die Daten müssen auf allen **REPLICAS** gespeichert werden.

□ **STUFEN DER LESEKONSISTENZ:**

- **ANY:** Nicht verfügbar.
- **ONE:** Liefert die Daten vom naheliegensten Knoten zurück
- **QUORUM:** Liefert die neuesten Daten auf Basis der Zeitstempel durch eine Ausführung eines **QUORUM READ**
- **ALL:** Liefert die neuesten Daten auf Basis der Zeitstempel aus allen **REPLICAS** zurück. □

QUELLEN:

- Architecture Overview:
<http://wiki.apache.org/cassandra/ArchitectureOverview>
- Internal Architecture:
<http://wiki.apache.org/cassandra/ArchitectureInternals>
- Data Model:
<http://wiki.apache.org/cassandra/DataModel>
- CQL:
<http://www.datastax.com/0.8/api/cql.ref>

COUCHDB

COUCHDB ist ein Akronym und steht für **CLUSTER OF UNRELIABLE COMMODITY HARDWARE DATABASE**.

DATENMODELL Die Speicherung der Daten erfolgt im **JSON** Format. COUCHDB baut auf einem DOKUMENTORIENTIERTEN DATENMODELL auf. Dokumente werden in COUCHDB bei einer Änderung nicht überschrieben, sondern es wird ein neues Dokument mit der selben ID generiert. Dokumente werden zur effizienteren Verarbeitung in **B-BÄUMEN** gespeichert. □

KOMMUNIKATIONSMODELL Die Kommunikation mit COUCHDB erfolgt über **HTTP**. Als Schnittstelle wird **REST** eingesetzt. Da es sich bei **HTTP** um ein zustandsloses Protokoll handelt, müssen alle relevanten Transaktionsdaten in einer Anfrage übergeben werden. Innerhalb einer Anfrage erfüllt CouchDB die **ACID** Kriterien. □

KONSISTENZMODELL **COUCHDB** wurde konzipiert, um eine hohe **DATENVERFÜGBARKEIT** aufzuweisen. Das wird dadurch erreicht, daß die Daten auch während **REPLIKATIONSVORGÄNGEN** weiter lesbar bleiben. Um ein **LOCKING** der Daten zu verhindern setzt COUCHDB auf MVCC. Durch diese Form der Konfliktauflösung gehen keine Daten verloren, jedoch kann nicht sichergestellt werden, daß die Daten in der aktuellen Version vorliegen. □

REPLIKATION COUCHDB ist als verteilte Datenbank konzipiert. Um die Konsistenz der Daten sicherzustellen, setzt COUCHDB auf **REPLIKATION**.

Die **REPLIKATION** der Daten findet als inkrementeller Prozess statt. Die einzelnen Knoten kommunizieren untereinander ebenfalls über die **REST** Schnittstelle. Hinzu kommt, daß das dokumentenweise Übertragen der Daten gegenüber Hardwareausfällen sicher ist. Die Architektur der Datenbank macht sie generell gegenüber Knotenausfällen robust. □

DATENMANIPULATION Als Abfragesprache wird in COUCHDB auf in JAVASCRIPT programmierte MAP/REDUCE Funktionen gesetzt. □

QUELLEN:

- Apache Projekt:
<http://couchdb.apache.org/>
- CouchDB Wiki:
<http://wiki.apache.org/couchdb/FrontPage>

HADOOP-HIVE

HADOOP HADOOP ist ein freies JAVA basiertes quelloffenes Framework, das unter der Apache Software Lizenz entwickelt wird. HADOOP ist eine Implementierung des MAP/REDUCE Algorithmus. Die Stärke von HADOOP besteht darin große Datenmengen in relativ geringer Zeit verarbeiten zu können.

Im Kern besteht HADOOP aus folgenden Komponenten:

- **HADOOP DISTRIBUTED FILE SYSTEM:** Ist das verteilte Dateisystem von Hadoop. Es zeichnet sich durch hohe Ausfallsicherheit aus. Mit Hilfe von **HDFS** werden die Daten auf die verschiedenen Knoten des **CLUSTERS** erteilt.
- **MAP/REDUCE FRAMEWORK:** Wird zur Datenverarbeitung eingesetzt. □

HIVE HIVE ist eine **DATA WAREHOUSE INFRASTRUKTUR** die auf HADOOP aufsetzt.

HIVE wird eingesetzt, um **ANALYSEN** und **ANFRAGEN** auf großen Datenbeständen auszuführen. Hive besitzt mit **HQL** eine **SQL** ähnliche Abfragesprache zur Interaktion mit dem Anwender.

Hive setzt in erster Line auf

- **SKALIERBARKEIT:** Hinzufügen von Knoten zum Hadoop Cluster
- **ERWEITERBARKEIT:** Hinzufügen von eigenen MAP/REDUCE Skripten
- **AUSFALLSTOLERANZ**

UPDATES, **TRANSAKTIONEN** und **INDIZES** werden von HIVE nicht unterstützt. HIVE greift zur Datenverarbeitung in erster Line auf MAP/REDUCE Skripten zurück. Updates werden über das Anlegen neuer Tabellen im Kontext der Ausführung von Skripten verwirklicht.

In traditionellen Datenbanksystemen wird beim Einspielen von Daten ein Schema erzwungen. Passen die Daten nicht zum Schema, werden sie verworfen. Dieser Ansatz wird auch **SCHEMA ON WRITE** genannt da die Daten beim Schreiben gegen das Schema geprüft werden. Hive hingegen arbeitet nach dem **SCHEMA ON READ** Ansatz. Dieser sieht vor, daß die Daten nicht schon beim Laden, sondern erst bei Anfragen gegen das Schema geprüft werden. □

FEATURES	SQL	HIVE
UPDATES	update insert delete	insert, overwrite, keine updates
TRANSAKTIONEN	ja	nein
INDIZES	ja	nein
LATENZ	Millisekunden	Minuten
MULTI-INSERTS	nein	ja
SELECT	sql 92	nur 1 Tabelle in der from Klausel
JOIN	sql 92	inner, outer, semi und map joins
UNTERABFRAGEN	sql 92	nur in der from klausel

QUELLEN:

- Apache Projekt:
<http://hive.apache.org/docs/r0.9.0/index.html>

MONGODB

MONGODB ist ein weiterer bekannter Vertreter für **NoSQL** Datenbanken. Namensgebend für MONGODB ist ein Wortspiel mit dem englischen **HUMONGOUS**.

EIGENSCHAFTEN

- Der MONGODB SERVER verwaltet einen Netzwerk von mehreren Datenbankinstanzen. Jeder Knoten kann einzeln **BEHANDELT** und **KONFIGURIERT** werden.
- **SKALIERBAR**
- **HOCHPERFORMANT**
- **SCHEMAFREI**
- **DOKUMENTORIENTIERT: BSON**
- **INTEGRIERTE QUERYLANGUAGE**
- **MONGODB** skaliert horizontal. Erreicht wird das durch den Einsatz von **REPLIKATION** und **SHARDING**.
- Schwächen zeigt die Datenbank in der Sicherheit.
- Ebenfalls sind **QUERIES** und **TRANSAKTIONEN** im Vergleich zu **RDBMS** nur verschlankt einsetzbar.
- **MAP/REDUCE** wird eingesetzt um Funktionen auf den Daten umzusetzen. **MAP/REDUCE** entspricht dem **GROUP OPERATOR**. □

QUELLEN

- Projekt:
<http://http://www.mongodb.org>
- BSON Spezifikation
<http://www.bsonspec.org>

APPEENGINE

BESCHREIBUNG Die APPEENGINE setzt auf BIGDATA zur

- **VERWALTUNG SEINER DATEN**
- **VERWALTUNG SEINER STRUKTUR**
- **SICHERSTELLUNG DER DATENKONSISTENZ** □

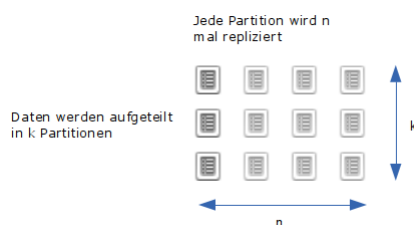
BIGTABLE BIGTABLE ist ein verteiltes Datenverarbeitungssystem das vom Konzept her als eine Tabelle ausgelegt ist, eine Tabelle die Petabytes an Daten speichert und über tausende Netzwerkknoten verteilt sein kann. BIGTABLE ist in der Lage, tausende von Anfragen in einer Sekunde zu verarbeiten.

BIGTABLE nimmt im **NoSQL** Bereich eine fundamentale Vorreiterrolle ein. Viele der entscheidenden Konzepte im **NoSQL** Bereich wurden für BIGTABLE pioniert. □

EIGENSCHAFTEN VON BIGTABLE

- **FUNDAMENTALLY DISTRIBUTED:** BIG TABLE wurde von Beginn an als verteilte, robuste Architektur entworfen. Die zu speichernden Daten sind verteilt auf eine Menge von Netzwerkknoten. Die Architektur ist so ausgelegt daß Knotenausfälle durch **REPLIKATION** kompensiert werden können.

DATA PARTITIONING wird eingesetzt um die Daten ihrem Schlüssel entsprechend auf verschiedenen Knoten zu verteilen. **REPLIKATION** sorgt auf der anderen Seite dafür daß Kopien von Daten auf mehreren Knoten verteilt vorliegen.



- **COLUMN ORIENTED:** BIGTABLE implementiert ein SPALTENORIENTIERTES DATEN-MODELL.
- **SEQUENTIAL WRITE:** BIGTABLE ist optimiert für Schreiboperationen.
- **CHUBBY: LOCK/FILE/NAME SERVICE.** Chubby realisiert einen Lockingmechanismus für Bigtable. Chubby wird in Bigtable für mehrere Aufgaben eingesetzt:
 - Sicherstellung daß es zur Ausführungszeit nur einen aktiven Master gibt
 - Speicherung von BIGTABLE Schema Informationen
- **MAP/REDUCE**
- **MVCC**
- **CONSISTENT HASHING** □

QUELLEN:

- Projekt:
<https://developers.google.com/appengine/>
- AppEngine Docs
<https://developers.google.com/appengine/docs/whatisgoogleappengine>

RDBMS und andere Datenbanksysteme

RELATIONALE DATENBANKSYSTEME

BESCHREIBUNG **RELATIONALE DATENBANKSYSTEME** wurden entwickelt um den wachsenden Ansprüchen einer neuen Generation von Applikationen gerecht zu werden.

- **REDUNDANZ UND INKONSISTENZ** von Daten
- **BESCHRÄNKTE ZUGRIFFSMÖGLICHKEITEN**
- **HERAUSFORDERUNGEN DES MEHRBENUTZERSBETRIEBS**
- **VERLUST VON DATEN**
- **INTEGRITÄTSVERLETZUNG**
- **SICHERHEITSPROBLEME** □

RELATIONALE DATENBANKEN: PRINZIPIEN

RELATIONALES DATENMODELL: Anfang der siebziger Jahre wurde das **RELATIONALE DATENMODELL** konzipiert. Die Besonderheit dieses Datenmodells besteht in der mengenorientierten Verarbeitung der Daten im Gegensatz zu den bis dahin vorherrschenden Datenmodellen.

Das **RELATIONALE DATENMODELL** ist im Vergleich zu anderen Modellen sehr einfach strukturiert. Es gibt im wesentlichen nur flache Tabellen (**RELATIONEN**), in denen die Zeilen den Datenobjekten entsprechen. In dieser einfachen Struktur liegt vermutlich der Erfolg der relationalen Datenbanken. Die in den Tabellen gespeicherten Daten werden durch entsprechende Operatoren ausschließlich mengenorientiert **VERKNÜPFT** und **VERARBEITET**.

Das **ENTITY-RELATIONSHIP** Modell wird verwendet, um die Struktur der Datenbank und die Relationen der Tabellen untereinander zu beschreiben. Das **EER** besitzt zwei grundlegende Strukturierungskonzepte:

- **ENTITYTYPEN:** Entsprechen den zu modellierenden Daten
- **BEZIEHUNGSTYPEN:** Werden eingesetzt, um die Beziehungen zwischen den Daten abzubilden.

Der große Vorteil des **RELATIONALEN DATENMODELLS** liegt in seiner Eigenschaft, komplexe Beziehungen zwischen Daten einfach modellieren zu können. Unter Zuhilfenahme verschiedener **CONSTRAINTS** kann die logische Konsistenz der Daten stets gewährleistet werden. □

RELATIONALE ANFRAGESPRACHE: SQL Die **RELATIONALE ALGEBRA** und der **RELATIONENKALKÜL** bilden die theoretischen Grundlagen für die Anfragesprache **SQL**.

Anfragesprachen wie **SQL** sind im Allgemeinen deklarativ. Die Benutzer geben lediglich an, welche Daten sie interessieren, und nicht, wie die Auswertung der Daten vorgenommen wird. Die oft sehr komplexen, zur Festlegung der Auswertung nötigen Entscheidungen werden vom Anfrageoptimierer des Datenbanksystems übernommen.

Zusätzlich zur Manipulation von Tabellen beinhalten Anfragesprachen auch Möglichkeiten zur Definition von **INTEGRITÄTSBEDINGUNGEN** für die Daten, zur Vergabe von **ZUGRIFFSRECHTEN** und zur **TRANSAKTIONSKONTROLLE** □

ELEMENTE VON SQL

- **DATENTYPEN:** **RELATIONALE DATENBANKEN** zeichnen die zu speichernden Daten aus, indem sie ihnen Datentypen zuordnen
- **SCHEMADEFINITION:** Mit dem Wissen der Datentypen können die Daten in Tabellen abgebildet werden. Die zu einer Datenbank gehörenden Tabellendefinitionen werden als **SCHEMA** bezeichnet.
- **SQL-ANFRAGEN:** Die **SQL ANFRAGEN** stellen eine mengentheoretische Auswertung des Datenbestandes dar. Der Grad der Komplexität der Anfrage wird bestimmt durch die Anfragestellung:
 - Anfragen über mehrere Relationen hinweg

- Aggregatfunktionen und Gruppierungen
- Geschachtelte Anfragen
- Quantifizierte Anfragen
- Sichten □

TRANSAKTIONEN Unter einer **TRANSAKTION** versteht man die Bündelung mehrerer Datenbankoperationen, die in einem verteilten System ohne unerwünschte Einflüsse durch andere Transaktionen als Einheit fehlerfrei ausgeführt werden sollen.

Eine Transaktion wird definiert durch folgende Eigenschaften:

- **ATOMICITY:** Abgeschlossenheit - Jede Operation in einer Transaktion ist vollständig abgeschlossen (oder es gibt gar keine)
- **CONSISTENCY:** Konsistenz - Bei Transaktionsstart und -ende ist die Datenbank in einem konsistenten Zustand.
- **ISOLATION:** Abgrenzung - Die Transaktion verhält sich so, als wäre sie die einzige Operation die auf der Datenbank ausgeführt wird
- **DURABILITY:** Dauerhaftigkeit - Ist eine Transaktion einmal abgeschlossen, wird sie nicht mehr rückgängig gemacht. □

KONSISTENZMODELL: Eine der herausragenden Eigenschaften von **RELATIONALEN DATENBANKEN** ist die stets gegebene Konsistenz der Daten. Ist eine performantere Verarbeitung der Daten gewünscht, besteht die Möglichkeit, die durch das **KONSISTENZMODELL** gegebenen Einschränkungen herabzustufen. Eine Schlüsseltechnologie hier sind die Transaktionen. □

VERTEILTE DATENBANKEN: Grundprinzip verteilter Datenbanken ist die **VERTEILUNGSTRASPARENZ**. Ein verteiltes Datenbanksystem muß sich dem Anwender gegenüber genauso verhalten wie ein nicht verteiltes.

Eine verteilte Datenbank ist ein logisch zusammengehörender Datenbestand, der physisch auf mehrere Knoten verteilt ist und durch das verteilte Datenbanksystem administriert wird. □

RELATIONALE DATENBANKEN: IMPLEMENTIERUNGEN

MYSQL Features:

- **THREAD POOLING**
- **MULTICORE SCALABILITY**

- **MYSQL CLUSTER**
- **SEMI SYNCHRONOUS REPLICATION**
- **PARTITIONING** □

11 G Features:

- **DATABASE REPLAY**
- **PARTITIONING**
- **SCHEMA MANAGEMENT**
- **DATA WAREHOUSING AND OLAP**
- **PL/SQL: EFFICIENT CODING**
- **SECURITY**
- **ORACLE EXADATA SIMULATOR**
- **SQL PLAN MANAGEMENT**
- **MANAGEABILITY**
- **SQL ACCESS ADVISOR**
- **SQL OPERATIONS: PIVOT AND UNPIVOT**
- **BACKUP AND RECOVERY**
- **RESILIENCY**
- **AUTOMATIC STORAGE MANAGEMENT**
- **COMPRESSION**
- **CACHING AND POOLING**
- **DATA GUARD**



Implementierung der Prototypen

EINFÜHRUNG

ÜBERSICHT Im Zuge der Arbeit haben wir einen Katalog von Anforderungen erarbeitet, der die Rahmenbedingungen zur Gestaltung der **ARCHITEKTUR** und technischen **UMSETZUNG** einer bestimmten Klasse von Applikationen definiert. Wir haben die unterschiedlichen **KONZEPTE** und **PARADIGMEN** kennengelernt, die für uns relevant sind, um das Problem zu verstehen und adäquat zu interpretieren.

Zur Etablierung unserer Ergebnisse werden 2 **PROTYPEN** erstellt und geprüft. Die **PROTYPEN** werden unterschiedlichen Paradigmen folgend entwickelt. Zum Vergleich werden die **PROTYPEN** identischen Lastentests unterzogen. Im Laufe dieses Kapitels wollen wir die unterschiedlichen Technologien vergleichen, die in die Gestaltung der Prototypen eingeflossen sind. □

RDBMS vs NoSQL Beide Prototypen implementieren das **MVC MODELL**. Die Herausforderung in der Gestaltung beider Programme ist es, sowohl das **VIEW** als auch den **CONTROLLER** für beide identisch zu halten. Lediglich in der **MODEL SCHICHT** streben wir eine grundlegende Differenzierung im Sinne der Implementierung an.

Die **PROTYPEN** weisen eine grundlegende Differenzierung in ihrem **DATEN-**, **KONSISTENZ-** und **REPLIKATIONSMODELL** auf. Ein **PROTOTYP** wird den Prinzipien der **NoSQL** Ansätze folgend gestaltet, der andere nutzt die Theorie der verteilten **RELATIONALEN DATENBANKEN** für die **DATENVERARBEITUNG**. □

PROGRAMMIERSPRACHE

VERTEILTE SYSTEME Die **GRUNDLEGENDE HERAUSFORDERUNG**, aber auch **HERAUSRAGENDE MÖGLICHKEIT** für heutige Softwareentwickler besteht darin verteilte Systeme zu **VERSTEHEN** und zu **ENTWICKELN**. Zeitgerechte Anwendungen, im **KLEINEN** aber auch im **INDUSTRIELLEN BEREICH**, werden mit hoher Wahrscheinlichkeit in einer Art von Netzwerk zur Anwendung kommen. Moderne **PROGRAMMIERSPRACHEN** werden im Licht dieser **ANFORDERUNGEN** entwickelt. □

PROGRAMMIERSPRACHEN - FUNKTIONALITÄT Viele Programmiersprachen integrieren bereits auf Sprachebene grundlegende Prinzipien zur Unterstützung der Entwicklung **VERTEILTER** und **PARALLELER** Software. Zusätzliche Funktionalität wird den Entwicklern mit **FRAMEWORKS** zur Verfügung gestellt.

Der Fortschritt und die Vielfalt webbasierter **CLIENT-SERVER** Frameworks erlauben uns aus einem Pool von Technologien zu wählen. Der Großteil dieser Technologien etabliert **HTML CLIENTS** mit Unterstützung für diverse Scriptsprachen, während am Server die zentrale Datenverarbeitung stattfindet.

Zur Realisierung unserer Prototypen wollen wir uns auf die Wahl einer Programmiersprache einschränken. Die Programmiersprache wird eingesetzt, um die erarbeitete Architektur umzusetzen.

Im Zuge dieses Abschnitts werden wir einige gängige Technologien **VORSTELLEN** und **VERGLEICHEN**. □

JAVA

JAVA ist eine **OBJEKTORIENTIERTE PROGRAMMIERSPRACHE**. Die Programmiersprache ist ein Bestandteil der **JAVA TECHNOLOGIE**. Die **JAVA TECHNOLOGIE** ist als Komposition zu verstehen aus einer **VIRTUELLEN MASCHINE**, **PROGRAMMIERSPRACHE**, **WERKZEUGEN** und **BIBLIOTHEKEN**.

JAVA TECHNOLOGIE (Quelle: [21]) JAVA setzt sich zusammen aus:

- **JDK:** JAVA-ENTWICKLUNGSWERKZEUGE zum Erstellen von Java Programmen
- **JRE:** JAVA-LAUFZEITUMGEBUNG zu deren Ausführung.
- **JVM:** LAUFZEITUMGEBUNG besteht selbst aus der **VIRTUELLEN MASCHINE** sowie den mitgelieferten **BIBLIOTHEKEN** der **JAVA LAUFZEITUMGEBUNG**. □

PROGRAMMIERSPRACHE JAVA Die Programmiersprache JAVA dient innerhalb der JAVA TECHNOLOGIE vor allem zum Formulieren von Programmen. Java Programme werden in 2 Schritten erstellt. Zunächst wird der **QUELLCODE** erstellt und im Anschluß daran der **PROGRAMMCODE** generiert.

- **QUELLCODE:** Dieser liegt zunächst als reiner, menschenverständlicher Text vor, als sogenannter **QUELLCODE**. Der Quellcode ist nicht direkt ausführbar
- **BYTECODE:** der Java-Compiler übersetzt den Quellcode in den **JAVA-BYTECODE**. Die JVM MASCHINE führt diesen Code aus.

Zweck dieser **VIRTUALISIERUNG** ist **PLATTFORMUNABHÄNGIGKEIT**. Das Programm soll ohne weitere Änderung auf jeder Rechnerarchitektur laufen können, wenn dort eine passende Laufzeitumgebung installiert ist. □

EIGENSCHAFTEN VON JAVA

- **EINFACHHEIT**
- **OBJEKTORIENTIERUNG**
- **VERTEILTHEIT**
- **ROBUSTHEIT**
- **SICHERHEIT**
- **ARCHITEKTURNEUTRALITÄT**
- **PORTABILITÄT**
- **PARALLELISIERBARKEIT**
- **DYNAMISCH** □

VERTEILTE SYSTEME JAVA JAVA wurde mit der Prämisse entwickelt eine Sprache zur Programmierung verteilter Systeme zu schaffen. Viele grundlegende Konzepte zur Programmierung verteilter Systeme sind direkt in die Sprache eingebettet:

- **MULTITHREADING**
- **SOCKETS, SERVERSOCKETS, URL**
- Unterstützung unterschiedlicher **KOMMUNIKATIONSPROTOKOLLE (TCP ,UDP)**
- **SICHERHEITSFEATURES** □

FRAMEWORKS Für JAVA steht eine immense Zahl an verschiedenen FRAMEWORKS zur Verfügung.

- **JSP**
- **JSF**
- **GRAILS**
- **EJB**
- **APACHE TAPESTRY**
- **GWT** □

EINSATZBEREICH Das Einsatzgebiet der JAVA PLATTFORM findet sich in erster Linie im ENTERPRISEBEREICH. JAVA erfährt eine kontinuierliche **WEITERENTWICKLUNG** und **EVOLUTION**. Die meisten namhaften Unternehmen setzen JAVA für die Entwicklung ihre Anwendungen ein.

- **TWITTER**
- **FOTOLOG**
- **AMAZON**
- **EBAY**
- **FLICKER**
- **AMAZON DYNAMO**
- **FEEDBURNER**
- **GOOGLESPEAK**
- **MAILINATOR** □

PHP

BESCHREIBUNG PHP ist eine Skriptsprache mit einer an **C** und **PERL** angelehnten Syntax, die hauptsächlich zur Erstellung dynamischer Webseiten oder Webanwendungen eingesetzt wird. PHP wird als freie Software unter der PHP-Lizenz verbreitet.

PHP zeichnet sich durch die breite **DATENBANKUNTERSTÜTZUNG** und **INTERNET-PROTOKOLLEINBINDUNG** sowie die Verfügbarkeit zahlreicher Funktionsbibliotheken aus. PHP wird auf etwa 244 Millionen Webseiten eingesetzt (Stand: Januar 2013).

PHP ist ein System, das PHP-Code serverseitig verarbeitet. Das bedeutet, dass der Quelltext nicht an den Webbrowser übermittelt wird, sondern an einen Interpreter auf dem Webserver. Erst die Ausgabe des PHP-Interpreters wird an den Browser geschickt. In den meisten Fällen ist das ein HTML-Dokument, wobei es mit PHP aber auch möglich ist, andere Dateitypen wie Bilder oder PDF-Dateien, zu generieren. □

WEBFRAMEWORKS

- CAPPUCCINO
- SEAGULL
- SYMFONY
- CAKEPHP □

ASP.NET

Das von Microsoft entwickelte .NET Framework ist ein bekannter JAVA Klone. Die für JAVA etablierten Ergebnisse sind in **BIJEKTIVER WEISE** auch für .NET gültig.

WAHL DER PROGRAMMIERSPRACHE: JAVA

BEGRÜNDUNG JAVA als Werkzeug ist weder besser noch schlechter geeignet zur Umsetzung **SKALIERENDER ARCHITEKTUREN** für **VERTEILTE APPLIKATIONEN IM DATEN-VERARBEITUNGSBEREICH** als andere Programmiersprachen. JAVA skaliert genauso wie jede andere Programmiersprache.

Was Java auszeichnet ist:

- **JAVA ALS PLATFORM:** Ein Ecosystem bestehend aus **JVM**, **BIBLIOTHEKEN** und **WERKZEUGEN**.
- **COMMUNITY:** Die Evolution von **JAVA** wird von einer zahlreichen und vielfältigen **COMMUNITY** vorangetrieben.
- **CLOUD COMPUTING STANDARD:** **ORACLE** arbeitet als eines der ersten Unternehmen an der Etablierung eines **STANDARDS** für eine einheitliche **CLOUD COMPUTING API**. Mit JAVA 8 soll diese teil des JDK werden. □

SKALIERENDE INFRASTRUKTUR Die Java Programmiersprache hat keinen Einfluß, ob ein Programm skaliert oder nicht. JAVA als Plattform dient zur Entwicklung neuer **SKALIERENDER TECHNOLOGIEN**. Der Architekturansatz, den Großteil der Arbeitslast auf einen einzelnen **MONOLITEN SERVER** auszulagern, wird zu Gunsten anderer komplexerer Technologien wie **JGROUPS**, **EHCACHE** und **TERRACOTTA** aufgegeben. JAVA als ECOSYSTEM versuchte seit seinem Bestehen die Entwicklung dieser Art von Software zu unterstützen. □

WAHL DES FRAMEWORKS: GWT-GAE

ANFORDERUNGEN Die BIBER WEBAPPLIKATION ist als **CLIENT-SERVER** System konzipiert. Wir werden versuchen, den gesammelten Anforderungen der Spezifikation mit der Wahl einer entsprechenden Technologie zu entsprechen. Die auszuwählende Technologie muß die Kriterien für ein **VERTEILTES SYSTEM** erfüllen:

- **HOHE VERFÜGBARKEIT**
- **ORTSTRANSparenZ**
- **REPLIKATIONSTRANSparenZ**
- **VERTEILTE ANFRAGEBEARBEITUNG**
- **VERTEILTE TRANSPARENZBEARBEITUNG**
- **HARDWAREUNABHÄNGIGKEIT**
- **BETRIEBSSYSTEMUNABHÄNGIGKEIT**
- **LEISTUNGSTRANSparenZ**

Gleichzeitig muß die Applikation als Anwendung für das WORLD WIDE WEB konzipiert sein. Alle WEBFRAMEWORKS scheinen in gleicher Weise für den Einsatz in unserem Prototypen geeignet zu sein. Als Technologien, die auf **HTTP** aufsetzen und eine **CLIENT-SERVER** Architektur etablieren, werden die geforderten Eigenschaften erfüllt. □

BEGRÜNDUNG GWT-GAE ist ein von GOOGLE entwickeltes Framework. GWT-GAE ist für GOOGLE CLOUD optimiert und stellt ein Webframework zur Entwicklung von **CLIENT SERVER** Architekturen bereit. Alle geforderten Anforderungen können im Gegensatz zu den anderen Frameworks mit diesem Werkzeug umgesetzt werden. Im Vergleich zu anderen Frameworks ist GWT-GAE **AUSGEREIFT** und **ERPROBT**. Zusätzlich stellt uns das **FRAMEWORK** die Werkzeuge zur Verfügung sowohl **NoSQL** als auch **RDBMS** Systeme für die Datenverarbeitung zu integrieren. □

VERGLEICH

- **BIBER WEBANWENDUNG:** Die primäre Anforderung an die Applikation ist es eine **PERFORMANTE, LINEAR SKALIERENDE** Architektur aufzubauen. GWT - GAE ist Teil der GOOGLE CLOUD. Damit ist das Framework geeignet eine solche Architektur zu implementieren.
- **RDBMS vs NoSQL:** GWT - GAE stellt sowohl eine Implementierung von BIGTABLE zur Verfügung, als auch eine Anbindung an CLOUDSQL mit MYSQL.
- **HTTP - CLIENT/SERVER:** GWT - GAE baut auf dem **HTTP PROTOKOLL** auf. □

NoSQL: IMPLEMENTIERUNG DES MODELS

NoSQL Systeme sind als eine Sammlung unterschiedlicher Konzepte und Technologien zu verstehen. Trotz diverser Technologieparadigmen gibt es einheitliche Ziele die jedes dieser Systeme umzusetzen versucht.

- **SCALE-OUT:** lineare Skalierung
- **ELASTICITY:** das System kann dynamisch zur Laufzeit erweitert werden

PARADIGMENVERGLEICH

Um die unterschiedlichen **NoSQL** Systeme vergleichen zu können erarbeiten wir uns einen Katalog von Kriterien um die Systeme im Anschluß entsprechend bewerten zu können:

- **READ/WRITE OPTIMIZED:** Das dem System zugrundeliegende Technologieparadigma resultiert oft in einem unterschiedlichen Verhalten für SCHREIB- und LESEZUGRIFFE
- **SYNC/ASYNC REPLICATION:** **REPLIKATION** wird eingesetzt um in einem verteilten System DATENVERLUST zu vermeiden, die GESCHWINDIGKEIT zu verbessern und für AVAILABILITY zu sorgen.

Wir unterscheiden zwei Formen der Replikation:

- **SYNCHRONE REPLIKATION:** SYNCHRONE REPLIKATIONEN stellt sicher daß die Daten in allen Knoten konsistent sind. Bei diesem Replikationsansatz nimmt man gleichzeitig hohe Latenzzeiten bei Datenupdates in Kauf.
- **ASYNCHRONE REPLIKATION:** ASYNCHRONE REPLIKATION hat keinen negativen Auswirkungen auf die Latenz des Systems, es kann aber vorkommen, daß in manchen Knoten alte Daten gespeichert werden.

SYSTEMVERGLEICH

Das **CAP THEOREM** (Kapitel 1, Seite.11) hilft uns eine **KLASSIFIZIERUNG** für **VERTEILTE SYSTEME** im **DATENVERARBEITUNGSBEREICH** zu etablieren. Im **NoSQL** Bereich beschränkt sich unsere Diskussion damit auf die CP- UND AP SYSTEME. Die folgende **SYSTEME** sollen für einen Vergleich näher betrachtet werden:

- **MONGODB**
- **HBASE**
- **BIGTABLE**

MONGODB vs BIGTABLE: VERGLEICH

- **REPLIKATION:** Beide unterstützen **SYNCHRONE REPLIKATION**
- **READ/WRITE OPTIMIZED:** Beide SYSTEME sind **SCHREIB/LESE** optimiert
- **DATENCHEMA:** Beide DB SYSTEME setzen auf ein **SCHEMA-LESS** Datenschema
- **ABFRAGEN:** Bei beiden Systemen muß bei der Gestaltung von **QUERIES** auf ihre Komplexität achtgegeben werden.
- **SKALIERBARKEIT:** Beide Systeme **SKALIEREN LINEAR**
- **TRANSAKTIONEN:**
 - **BIGTABLE** unterstützt **TRANSAKTIONEN** solange die entsprechenden Daten in Gruppen strukturiert wurden.
 - **MONGODB** unterstützt keine **TRANSAKTIONEN** als solche, es gibt aber **ATOMIC OPERATIONS** die helfen **TRANSAKTIONEN** zu simulieren.

MONGODB vs BIGTABLE: RESULTAT Beide DB SYSTEME bieten ähnliche Features, es muß aber erwähnt werden daß **GOOGLES BIGTABLE** das Ältere der beiden Systeme ist und mit der **IMPLEMENTIERUNG** von **MONGODB** nicht mithalten kann. **MONGODB** bleibt unerreicht solange ganze Familien von **OBJEKTGRAPHEN** gespeichert oder geladen werden sollen.

Für den Einsatz von **GOOGLES APPENGINE** spricht in letzter Instanz jedoch nur **GOOGLES** unerreichte **HARWARE KONFIGURATION** auf der die **APPENGINE** aufsetzt. □

HBASE vs BIGTABLE: VERGLEICH **APACHE HBASE** ist eine **OPEN-SOURCE NoSQL** Datenverarbeitungssystem modelliert nach **GOOGLES BIGTABLE**. Im Gegensatz zum **BIGTABLE**, das auf **GOOGLES FILE SYSTEM** aufsetzt, verwendet **HBASE HADOOP** als Kern.

Vergleich:

- **REPLIKATION:** **HBase** unterstützen **SYNCHRONE REPLIKATION**
- **READ/WRITE OPTIMIZED:** Der Zugriff auf Daten ist **STRENG KONSISTENT** jedoch nicht immer sehr performant unter **HBASE**
- **DATENCHEMA:** Beide DB SYSTEME setzen auf ein **SCHEMA-LESS** Datenschema

- **ABFRAGEN:** Das System ist in der Lage große Datenmengen in kürzester Zeit zu verarbeiten. HBASE setzt dazu **MAP/REDUCE** ein.
- **SKALIERBARKEIT:** linear skalierbar

HBASE VS BIGTABLE: RESULTAT Wieder fällt eine große Ähnlichkeit zwischen den beiden System auf. Wieder muß sich **BIGTABLE** gefallen lassen das ältere von beiden Systemen zu sein. **HBASE** wird jedoch aufgrund seines Kerns in erster Linie zum VERARBEITEN und ANALYSIEREN von Daten eingesetzt. Im Gegensatz dazu ist **BIGTABLE** optimiert für einfache SCHREIB-/LESEZUGRIFFE.

Erneut fällt die Wahl zugunsten von **BIGTABLE** aus da **HBASE** in erster Linie zur DATENAUSWERTUNG und -VERARBEITUNG eingesetzt wird. □

RDBMS - IMPLEMENTIERUNG DES MODELS

BESCHREIBUNG Verteilte **RELATIONALE DATENBANKSYSTEME** werden seit Jahren **EINGESETZT** und **WEITERENTWICKELT**. **RDBMS** Systeme finden Verwendung für kleine Anwendungen, genauso wie in Banken und anderen datenintensiven Umgebungen. □

RESULTAT Die den **RDBMS** Systemen zugrundeliegende Technologie erfährt nun seit Jahren eine Entwicklung und Verfeinerung. Die Produkte der verschiedenen RDBMS SYSTEM Anbieter weisen zwar im Detail Unterschiede auf, diese sind jedoch marginaler Natur und fallen kaum ins Gewicht. Für die **IMPLEMENTIERUNG** des **RDBMS** Models wählen wir **MYSQL** weil es sich um eine sehr gute OPEN SOURCE Datenbank handelt. □

Architekturbeschreibung

EINE PATTERNSPRACHE

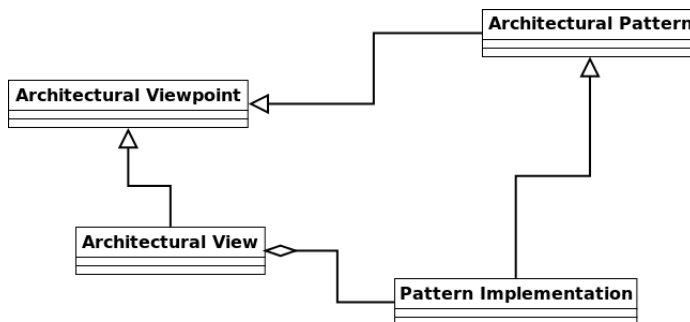
ARCHITECTURAL PATTERNS ARCHITECTURAL PATTERNS sind ein grundlegendes Werkzeug im Bereich der Softwareentwicklung. Sie stehen stellvertretend für einen Katalog diverser Alternativen zur Gestaltung umfangreicher Softwareanwendungen. Der Einsatz von ARCHITECTURAL PATTERNS hilft uns die Zahl der möglichen Architekturansätze auf ein überschaubares Maß einzuschränken. □

PATTERNSPRACHE Wir wollen uns des Konzeptes der PATTERNSPRACHE bedienen um eine strukturierte Klassifizierung der Designmöglichkeiten zu schaffen. Eine PATTERNSPRACHE versteht sich als Supermenge existierender Sammlungen von ARCHITECTURAL PATTERNS. Die Klassifizierung der Pattern baut auf dem Konzept der ARCHITECTURAL VIEWS auf.

Beschreibung der Elemente des Konzepts: (siehe [22])

- **ARCHITECTURAL VIEW:** Ist eine Repräsentation des Systems aus der Perspektive unterschiedlicher verwandter BELANGE. (Ein BELANG in einem verteilten Softwaresystem beschreibt die Abbildung der Softwarekomponenten auf Netzwerkknoten)
- **ARCHITECTURAL VIEWPOINT:** Die Typen der Elemente und deren Beziehungen werden durch die VIEWPOINTS beschrieben. Eine VIEW ist eine Instanz eines VIEWPOINTS für ein bestimmtes System.

Grundlegend für PATTERNSPRACHEN ist das Konzept der VIEWS. Views geben uns die Möglichkeit Pattern in Perspektive bestimmter Schwerpunkte zu betrachten. □



VIEWLISTE

- **LAYERED VIEW:** Das System wird dargestellt als ein komplexes System heterogener Komponenten, die sich als autonome interagierende Teile verstehen.
- **DATA FLOW VIEW:** Das System versteht sich als Folge von Transformationen auf einem Stream von Eingabedaten.
- **DATA CENTERED VIEW:** Das System wird gesehen als ein persistenter, gemeinsam genutzter Datenport, genutzt und verändert von diversen Komponenten.
- **ADAPTION VIEW:** Das System wird als Summe unveränderlicher und adaptionsbereiter Teile verstanden.
- **USER INTERACTION VIEW:** Das System wird primär aus der Perspektive des Anwenders dargestellt. Wir unterscheiden zwischen GUI Teilen und Komponenten, die die Applikationslogik implementieren.
- **COMPONENT INTERACTION VIEW:** Das System wird als eine Zahl unabhängiger Komponenten die miteinander zusammenarbeiten verstanden.
- **DISTRIBUTION VIEW:** Das System wird gesehen als Summe von Komponenten die auf ein Netzwerk verteilt werden. □

TOP LEVEL ARCHITEKTUR

LAYERS PATTERN

ARCHITEKTUR Zur Beschreibung der obersten Ebene unserer Architektur wählen wir das **LAYERS PATTERN**. Um eine Vergleichbarkeit der Prototypen zu gewährleisten, soll die Architektur beider Programme dem Prinzip des **SCHICHTENMODELLS** folgend gestaltet werden. Gelingt es uns einen Großteil der Schichten in beiden Prototypen identisch zu halten, kann ein legitimer Vergleich der Programme durchgeführt werden.

Das **MVC MODELL** ist eine Implementierung des **SCHICHTENMODELLS**, eingesetzt zur Strukturierung einer bestimmten Klasse von Anwendungen.

Die einzelnen Schichten des Modells sind

- **MODEL:** Das MODEL als Schicht ist für die Datenverarbeitung der Anwendung verantwortlich
- **VIEW:** Das VIEW dient in jeder Weise als Anwenderschnittstelle. Diese Schicht erlaubt den Zugriff des Anwenders auf die Funktionalität des Programmes.
- **CONTROLLER:** Der CONTROLLER dient als Mediator zwischen **MODEL** und **VIEW**.
□

DIFFERENZIERUNG Beide PROTOTYPEN implementieren das **MVC MODELL**.

Die Prototypen weisen eine grundlegende Differenzierung in ihrem

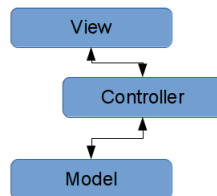
- **DATENMODELL**
- **KONSISTENZMODELL**
- **REPLIKATIONSMODELL** auf.

Der ein PROTOTYP wird den Prinzipien der **NoSQL** Ansätze folgend gestaltet, der andere nutzt die Theorie der verteilten **RELATIONALEN DATENBANKEN** für die Datenverarbeitung.

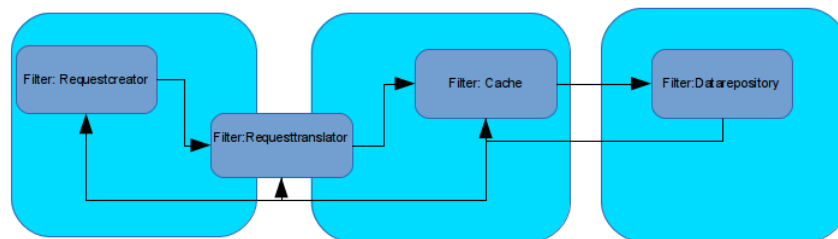
Das **MVC MODELL** ist ein etabliertes Pattern im Bereich der Webanwendungen. In diesem Zusammenhang scheint unsere Entscheidung für das **MVC PATTERN** nachvollziehbar und logisch. Das grundlegende Konzept des WORLD WIDE WEBS als CLIENT-SERVER System zwingt uns eine entsprechende Modularisierung unserer Anwendung auf. Die Layerstruktur ist eingebettet in eine einschließende CLIENT-SERVER Architektur. Zur Entlastung des Servers setzen wir auf SMART CLIENTS und serverseitig auf REPLIKATION. □

DIAGRAMME

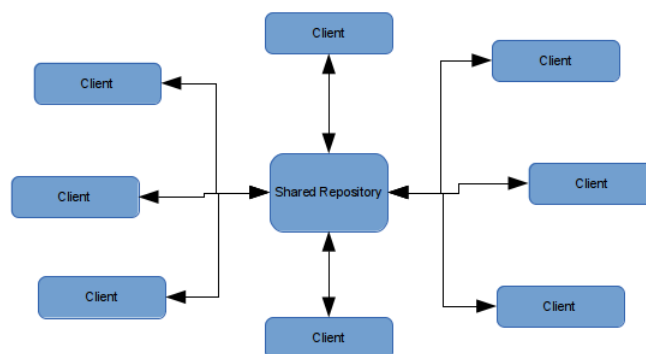
LAYERED VIEW Darstellung der obersten Architekturschicht in einem LAYERED VIEW



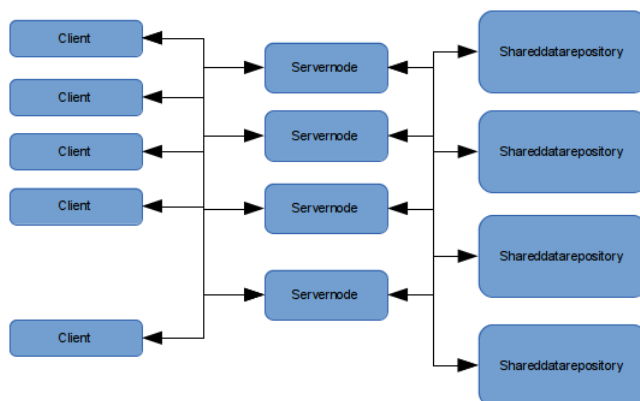
DATAFLOW VIEW Darstellung der obersten Architekturschicht in einem DATAFLOW VIEW



DATACENTERED VIEW Darstellung der obersten Architekturschicht in einem DATA-CENTERED VIEW



COMPONENTINTERACTION VIEW Darstellung der obersten Architekturschicht in einem COMPONENTINTERACTION VIEW



ALTERNATIVE DESIGNMÖGLICHKEITEN

Im Lichte eines wissenschaftlichen Diskurses möchten wir andere ARCHITECTURAL PATTERNS auf ihre Tauglichkeit für einen Einsatz in unserer Anwendung prüfen.

Folgender Katalog an möglichen Patterns soll näher diskutiert werden.

- **SOA**
- **ETL**
- **EAI/ESB**

DATA EXTRACTION TRANSFORMATION AND LOADING (ETL) Das ETL PATTERN findet in erster Line Einsatz in ENTERPRISE DATA WAREHOUSE Anwendungen. Das Pattern ist primär darauf ausgerichtet **DATENANALYSE** und **REPORTING** transparent und performant zu gestalten. Wir bemerken einen grundlegenden Schwerpunkt im **DATEN-VERARBEITUNGSBEREICH**.

Aufgaben eines ETL SYSTEMS:

- Das Pattern ist primär darauf ausgerichtet **DATENANALYSE** und **REPORTING** transparent und performant zu gestalten
- Ein ETL SYSTEM besteht aus unabhängigen Komponenten, die nach Belieben vernetzt werden
- ETL unterstützt parallele Datenverarbeitung immenser Datenmengen

Das primäre Augenmerk in der Gestaltung der Architektur des Prototypen dient dem **LAUFZEITVERHALTEN** der Applikation und der **LINEAREN SKALIERBARKEIT** der Anwendung, beides Kriterien, die im Zusammenhang des ETL PATTERN zweitrangig sind

oder keine Bedeutung finden. ETL SYSTEME sind kaum optimiert in Hinsicht auf ihr Laufzeitverhalten. Globales Transaktionshandling ist kritisch. □

EAI/ESB Ein ENTERPRISE SERVICE BUS ist ein **ARCHITEKTUR MODELL**, das eingesetzt wird um die **INTERAKTION** und **KOMMUNIKATION** von untereinander unabhängigen Softwareanwendungen zu gestalten. Als Softwarearchitektur für verteilte Systeme ist ESB eine Abstrahierung der **CLIENT/SERVER** Architektur, die in erster Linie auf **TRANSPARENZ** und **FLEXIBILITÄT** im Zusammenhang der Kommunikation und Interaktion setzt. ESB wird primär eingesetzt in ENTERPRISE APPLICATION INTEGRATION (EAI).

Die Aufgaben eines **ESB** sind:

- **MONITORING** des Nachrichtenaustausches unter Applikationen
- **VERWALTUNG** des Versioning von Applikationen
- **NACHRICHTEN** und **EREIGNIS QUEUEING**
- **SICHERHEITSMANAGEMENT**
- **AUSNAHMEMANAGEMENT**

ESB wird eingesetzt, um **INTERAKTION** und **KOMMUNIKATION** von Services zu steuern. Im Zusammenhang unserer Anwendung besteht jedoch keine Notwendigkeit, mehrere Services untereinander zu vernetzen. Einige von einem ESB zur Verfügung gestellten Features finden auch in unserer Anwendung Eingang, unsere primären Anforderungen liegen jedoch in anderen Bereichen. □

SERVICE-ORIENTED ARCHITECTURE (SOA) SOA ist ein Softwaredesignansatz, der komplexe Softwareanwendungen als Komposition mehrerer unabhängiger Services versteht. Jedes Service stellt eine klar definierten Katalog an Funktionalität zur Verfügung. Ein Service wird unabhängig von den anderen Services erstellt.

SOA wird eingesetzt, um die Services untereinander zu vernetzen und eine homogene Funktionalität für den User zu erzeugen. Die Services können in unterschiedlichen Programmiersprachen erstellt werden. SOA integriert die unterschiedlichen Services zu einem Ganzen.

Eigenschaften von SOA:

- **STANDARDIZED SERVICE CONTRACT**
- **SERVICE LOOSE COUPLING**
- **SERVICE ABSTRACTION**
- **SERVICE REUSABILITY**

- **SERVICE AUTONOMY**
- **SERVICE STATELESSNESS**
- **SERVICE DISCOVERABILITY**
- **SERVICE COMPOSABILITY**

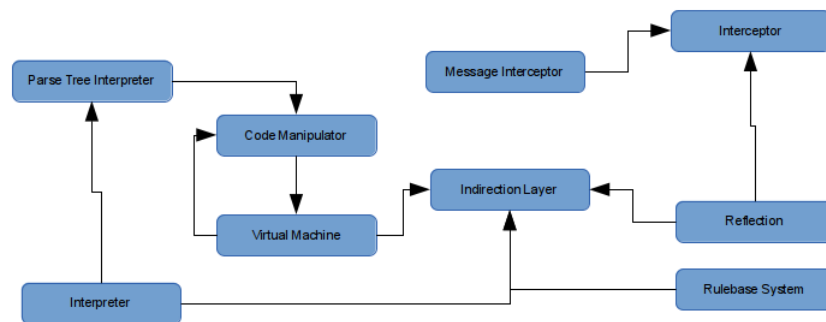
SOA wird eingesetzt, um Interaktion und Kommunikation von Services zu steuern. Im Zusammenhang unserer Anwendung besteht jedoch keine Notwendigkeit, mehrere Services untereinander zu vernetzen. Einige von SOA zur Verfügung gestellten Features finden auch in unserer Anwendung Eingang, unsere primären Anforderungen liegen jedoch in anderen Bereichen. □

VIEW

Das VIEW dient dem Anwender als Schnittstelle zur Anwendung. Die Architektur der VIEW Schicht wird vorgegeben durch die von der Anwendung vorgegebenen Anforderungen.

WEBBROWSER WEBBROWSER stellen die **PRIMÄRE CLIENTSEITIGE TECHNOLOGIE-KOMPONENTE** im WORLD WIDE WEB dar. Die Architektur eines Browsers entspricht der eines INTERPRETERS. Die Aufgabe des Browsers ist es HTML DOKUMENTE zu interpretieren und infolge zur Interaktion mit dem Anwender in einer grafischen Schnittstelle zur Verfügung zu stellen. In diesem Sinne wird uns die **TECHNOLOGISCHE** als auch **ARCHITEKTURELLE** Entscheidung auf oberster Ebene genommen. □

ARCHITEKTUR



CONTROLLER

Der CONTROLLER als Schicht dient als Mediator zwischen der **VIEW-** und der **MODELSCHICHT**. Der CONTROLLER wird sowohl auf der **CLIENT KOMPONENTE** als auch auf der **SERVER KOMPONENTE** deployed.

ARCHITEKTUR

Um die Kommunikation zu den 2 anderen Schichten zu etablieren, wollen wir den CONTROLLER in 4 Teile zerlegen:

- **MODEL**
- **VIEW**
- **PRESENTER**
- **APPCONTROLLER**

MODEL Dieser Teil des CONTROLLERS fungiert als Interface zum MODEL. Wir finden folgende Prinzipien:

- **ENTITYKLASSEN:** Kodierung der unterschiedlichen Entitäten
- **ENTITYDETAIL:** Eine abgespeckte Version der Entityklassen. Die Detailklassen werden verwendet, um Daten über das Netzwerk zu schicken. □

VIEW Dieser Teil des CONTROLLERS fungiert als Interface zur VIEWSCHICHT. Im VIEW befinden sich alle **UI KOMPONENTEN**, aus denen unsere Anwendung besteht. Die VIEWS sind ebenfalls verantwortlich für das Layout der UI Komponenten. Die VIEWS als solche sind nicht gebunden an die Entitäten. Die Information, die verwaltet werden soll, bezieht sich ausschließlich auf UI Komponenten.

Für das Wechseln zwischen verschiedenen VIEWS setzten wir auf HISTORY MANAGEMENT □

PRESENTER Der PRESENTER implementiert die Logik der CONTROLLERSCHICHT. Er beinhaltet eine Brücke zum

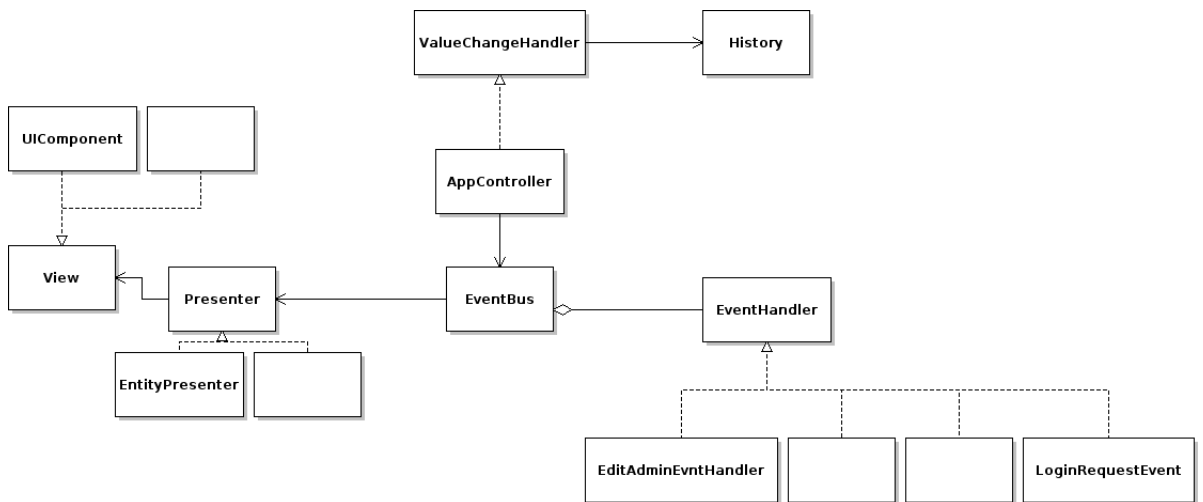
- **HISTORYMANAGEMENT**
- **VIEWTRANSITION**
- **ASYNCHRONES RPC**

Als Faustregel gibt es für jede VIEW einen PRESENTER der verantwortlich für die Verwaltung der VIEW und der unterschiedlichen Events der UI Komponenten ist. □

APPCONTROLLER Um Logik zu implementieren die nicht spezifisch ist für die PRESENTER, setzen wir auf die APPCONTROLLER Komponente. Das **HISTORYMANAGEMENT** und die **VIEWTRANSITION** Logik werden ebenfalls hier umgesetzt.

Der APPCONTROLLER ist zusätzlich verantwortlich für die **INITIALISIERUNG** und **VERWALTUNG** des EVENTBUSES und der Events. □

EVENTBUS Haben die PRESENTER ersteinmal einen EVENT der UI Komponenten abgesetzt, müssen wir in der Lage sein, auf jene zu reagieren. Wir setzen auf einen EVENTBUS, um die Events zu verwalten und an registrierte Komponenten weiterzuschicken. □

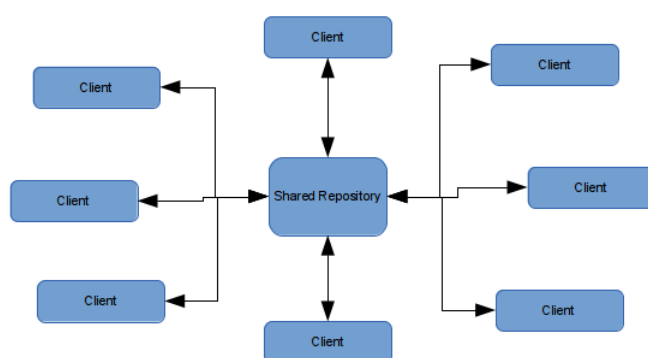


MODEL

Um eine Vergleichbarkeit der beiden Prototypen etablieren zu können, versuchen wir die beiden Anwendungen so ähnlich, wie möglich zu halten. Von der Architektur her entsprechen sich die MODELSCHICHTEN beider PROTOTYPEN, hinter den **ABSTRAKTEN KONZEPTEN** stehen jedoch unterschiedliche IMPLEMENTIERUNGEN.

ARCHITEKTUR

DATACENTERED VIEW Darstellung des MODELS im DATACENTERED VIEW



COMPONENTINTERACTION VIEW Die grundlegenden 2 Komponenten der MODELSCHICHT sind:

- **ENTITYMANAGER:** Der ENTITYMANAGER wird eingesetzt als **FASSADE**. Hinter der Implementierung des ENTITYMANAGERS stehen die unterschiedlichen Datenverarbeitungsbezogenen Technologien: **RDBMS** und **NoSQL**.

Die FASSADE dient uns als Interface zu den unterschiedlichen Technologien.

- **ENTITYKLASSEN**



Durchführung der Lastentests

EINFÜHRUNG

LASTENTESTS werden eingesetzt, um das **LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** der Prototypen zu bewerten und zu verstehen. Das LAUFZEITVERHALTEN ist in diesem Sinne als die Summe der unterschiedlichen Paradigmen für **DATENKONSISTENZ**, **REPLIKATION** und des **DATENMODELLS** zu beurteilen.

Beide Prototypen werden denselben LASTENTESTS unterzogen. Der Vergleich der Ergebnisse der Lastentests gibt Auskunft über die wahre Beschaffenheit der unterschiedlichen Konzepte. Die gefundenen Ergebnisse helfen uns, die gewünschten Eigenschaften zu prüfen und eine finale Architekturentscheidung zu treffen. □

TEST FRAMEWORK

Zur Durchführung der **LASTENTESTS** wird ein **TESTCLIENT** erstellt. Der **TESTCLIENT** wird eingesetzt um **HTTP REQUESTS** an den Server zu schicken. Am Server werden als Antwort darauf mehrer Operationen ausgelöst um die Datenverarbeitungskomponente zu testen. □

TESTGENERATOR ARCHITEKTUR

Der **TESTCLIENT** ist ein Java Programm zur Ausführung von **TESTCASES** und zur Evaluierung der Prototypen. Die primären Komponenten des **TESTCLIENTS** sind:

- **WORKEXECUTOR:** Der **WORKEXECUTOR** führt **WORKLOADS** aus. Jedem **WORKLOAD** werden mehrer **THREADS** zugewiesen um die anstehenden Operationen auszuführen. Diese Komponente steuert ebenfalls den **WORKLOAD** und **THROUGHPUT** der **TESTCASES**.

Der **WORKEXECUTOR** ist im allgemeinen zuständig für die Generierung von Daten und dem Ausführen von Operationen.

- **THREADPOOL:** Der **THREADPOOL** ist eine ressourceeffiziente Struktur zur Verwaltung der **THREADS**.
- **INTERFACE LAYER:** Der **INTERFACE LAYER** stellt die Schnittstell des Programms zu den Prototypen dar. □

WORKLOAD

Ein **WORKLOAD** bestimmt die Zusammensetzung von **SCHREIB-/ LESEOPERATIONEN**, **DATENGRÖSSE** und **REQUESTVERTEILUNG** eines Testcases. Als **SCHREIBOPERATION** verstehen wir alle **CRUD OPERATIONEN** ohne den **LESEOPERATIONEN**.

Wir sind in der Lage unterchiedliche Arten von Testdaten zu generieren:

- **READ-HEAVY WORKLOAD:** Der **WORKLOAD** beinhaltet 5% **UPDATE** und 95% **READ** Operationen.
- **WRITE-HEAVY WORKLOAD:** Der **WORKLOAD** beinhaltet 95% **UPDATE** und 5% **READ** Operationen.
- **BALANCED WORKLOAD:** Der **WORKLOAD** beinhaltet 50% **UPDATE** und 50% **READ** Operationen. □

TESTKRITERIEN

Die Tests sind ausgelegt um mehrere unterschiedliche Kriterien im System zu evaluieren. Wir versuchen die folgenden Größen zu testen:

- **PERFORMANCE**
- **SKALIERBARKEIT**
- **ELASTISCHE VERHALTEN**
- **SCHREIB-/LESEVERHALTEN**

PERFORMANCE Die PERFORMANCE bezogenen TESTCASES werden durchgeführt um die durch datenbezogene Operationen hervorgerufenen Latenzzeiten im System zu evaluieren.

Durch kontinuierliche Steigerung der Systemlast soll, bei gleichzeitiger Messung der Latenzzeiten jener Punkt gefunden werden, an dem das System an seine Grenzen stößt. Im Rahmen dieser Tests wird bewußt auf Skalierung verzichtet. □

SKALIERUNG Die SKALIERUNG bezogenen TESTCASES werden eingesetzt um die Geschwindigkeit des Systems zu messen im Bezug auf die Eigenschaft des System zu skalieren. In diesem Zusammenhang evaluieren wir 2 Größen:

- **SCALE UP FRAGESTELLUNG:** Wie verhält sich das System während die Zahl der Knoten erhöht wird?

Um die Tests durchzuführen, evaluieren wir die Performance des Systems in einem bestimmten Ausgangszustand. Anschließend werden die Daten gelöscht, die Zahl der Knoten erhöht und die Tests erneut durchgeführt. Skaliert das System sollte die **PERFORMANCE** gleich bleiben.

- **ELASTIC SPEEDUP FRAGESTELLUNG:** Wie verhält sich das System während die Zahl der Knoten zur Laufzeit erhöht wird?

Die Tests werden durchgeführt und ausgewertet während schrittweise die Zahl der Knoten erhöht wird. □

TEST-SETUP

Zum Durchführen der **TESTS** werden die **DATENVERWALTUNGSSYSTEME** zunächst mit Testdaten versorgt. Das **TESTFRAMEWORK** ist darauf ausgelegt die entsprechenden Testdaten automatisch zu generieren und den entsprechenden **WORKLOADS** zuzuordnen. Im folgenden soll die Beschaffenheit der **TESTDATEN**, des **WORKLOADS** und der **HARDWAREKONFIGURATION** des Systems beschrieben werden.

TESTDATEN Jeder **EINTRAG** hat eine Größe von 500 BYTES. Ein **EINTRAG** setzt sich in diesem Sinne zusammen aus

- **SCHLÜSSEL**: Typ: VARCHAR
- **FELDER**: Typ: VARCHAR □

WORKLOAD Jeder **WORKLOAD** ist so angelegt das 200 000 000 Einträge verarbeitet werden. Zusammenstellung des **WORKLOADS**:

- 200 000 000 **EINTRÄGE**
- Größe des Eintrags: 500 Byte
- 5 Felder zu 100 Bytes pro **EINTRAG**
- Parallele Verarbeitung der Tests (**THREADPOOL**) □

GOOGLE CLOUD Beide **PROTOTYPEN** werden als Teil der **GOOGLE CLOUD** deployed. Das System stellt uns sowohl eine Anbindung an **GOOGLES BIGTABLE**, als auch an ein **RDBMS(MYSQL)** System bereit. Die Test werden in der **GOOGLE CLOUD** ausgeführt. (Stand der Technologie **MAI 2013**).

Die **GOOGLE CLOUD** kommt automatisch mit folgenden Features:

- **LINEARE SKALIERUNG**
- **FAST DATA ACCESS**
- **UNLIMITED STORAGE**
- **BIGQUERY**

Je nach Kontrakt werden dem Kunden unterschiedliche **HARDWARE KONFIGURATIONEN** zur Verfügung gestellt.

Die Test wurden auf **SERVERCLUSTERN** mit folgender **HARDWARE** ausgeführt:

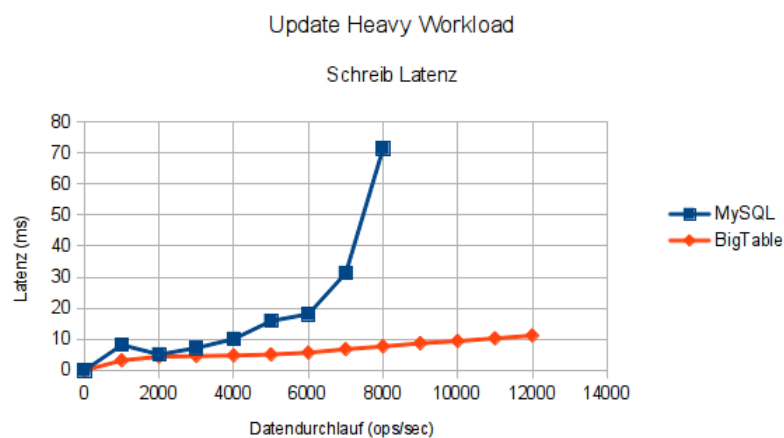
- **PROZESSOR**: 64 Bit quad core CPU

- **RAM:** 3.75 GB
- **HD:** nach Vereinbarung

TESTDURCHFÜHRUNG

PERFORMANCE-TEST

SCHREIBOPERATIONEN Um die Beschaffenheit des **SCHREIBVERHALTENS** der beiden **Prototypen** zu testen wird das **UPDATE-HEAVY WORKLOAD** geladen und bei kontinuierlicher Steigerung der Last die **LATENZ** der Systeme gemessen.

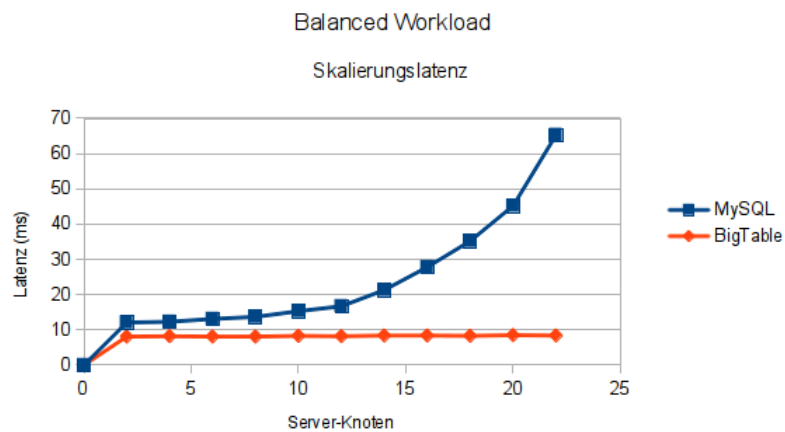


LESEOPERATIONEN Um die Beschaffenheit des **LESEVERHALTENS** der beiden **Prototypen** zu testen wird das **READ-HEAVY WORKLOAD** geladen und bei kontinuierlicher Steigerung der Last die **LATENZ** der Systeme gemessen.



SKALIERBARKEITS-TEST

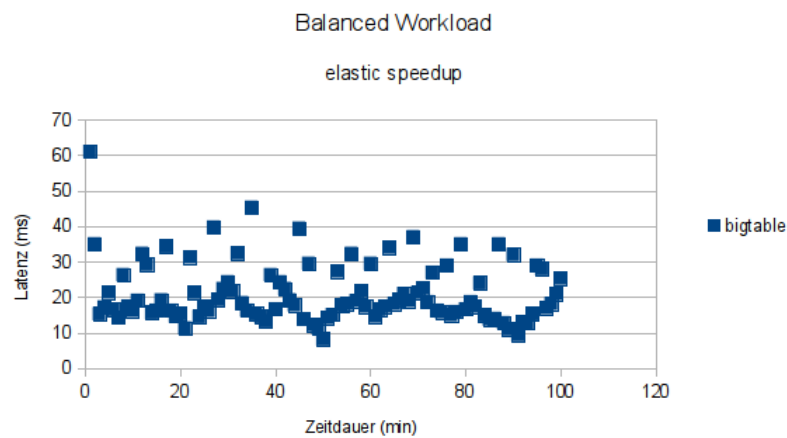
Um die **SKALIERUNGSFÄHIGKEIT** der beiden PROTOTYPEN zu messen wird das **BALANCED WORKLOAD** geladen. Im Gegensatz zu den PERFORMANCETESTS, bei denen die Zahl der SERVERKOTEN während des Tests **KONSTANT** gehalten wurde, wird nun die Zahl der SERVER **PROPORTIONAL** mit der Zahl der OPERATIONEN erhöht.



ELASTIC SPEEDUP

Ein VERTEILTES SYSTEM das in der Lage ist **SERVER KNOTEN**, ohne erwähnenswerte **LATENZ**, dem System hinzuzufügen besitzt die Eigenschaft des **ELASTIC SPEEDUP**. Wir laden das **BALANCED WORKLOAD**. Anschließend werden **SERVER KNOTEN** dem System hinzugefügt. In diesem Fall messen wir die Zeit die das System braucht um den neuen **SERVER KNOTEN** zur Laufzeit in das System zu integrieren. Ein SERVER nach dem anderen wird dem SYSTEM einzeln hinzugefügt. Anschließend wird gewartet bis sich das System wieder stabilisiert hat und keine zusätzliche **LATENZ** mehr erzeugt.

MYSQL für sich erfüllt in keiner annehmbaren Zeitdauer die Kriterien für **ELASTICHES SPEEDUP**. Die Messungen stabilisieren sich erst nach einer geraumen Zeit.



RESULTAT

PERFORMANCE-TESTS

Die **PERFORMANCE TESTS** sollen helfen die grundlegenden Eigenschaften (**CRUD OPERATIONEN**) der beiden **PROTOTYPEN** zu messen. Die **PERFORMANCE-TESTS** werden ohne **REPLIKATION** durchgeführt. □

READHEAVY WORKLOAD Um die Fähigkeit der **SYSTEME** zu messen, große Mengen von Daten schnell lesen zu können, wird ein **WORKLOAD** geladen der sich zu 95% aus **LESEOPERATIONEN** zusammensetzt. Bei der Durchführung des Test wird die Zahl der Operationen kontinuierlich erhöht und die **LATENZ** der **SYSTEME** gemessen. Beide **PROTOTYPEN** erzielen bei diesem **TEST** sehr gute Ergebnisse. Die gemessene **LATENZ** bleibt mit steigender **SYSTEMLAST** über die längste Zeit hinweg nahezu **LINEAR**. Erst ab einem gewissen **WERT** für den **DATENDURCHLAUF** erkennen wir einen signifikanten Anstieg in der **LATENZ**.

Daß **MYSQL** in dieser Testreihe leicht zu favorisieren ist, ist auf den Umstand zurückzuführen, daß **RDBMS** Systeme im Allgemeinen **LESE OPTIMIERT** sind. **RDBMS** Systeme erlauben es nicht nur der Anwendung schnell, große Mengen von Daten lesen zu können, es wird zusätzlich Funktionalität zur Verfügung gestellt um Daten schnell abfragen zu können. Dazu werden die **DATEN** im System indiziert. Was für das **LESEVERHALTEN** in diesem Fall günstig ist, erweist sich für das **SCHREIBVERHALTEN** als nachteilhaft. Werden neue **DATEN** geschrieben muß der **INDEX** jedesmal aufs neue konfiguriert werden.

Besteht der gegebene **WORKLOAD** eines **DATENVERARBEITUNGSSYSTEMS** in erster Linie aus **LESEOPERATIONEN** oder grundlegenden **AGGREGATFUNKTIONEN** (z.B.: **SUM()**, **AVG()**) so ist in diesem Fall den **RDBMS** Systemen ein Vorzug gegenüber den **NoSQL** Systemen zu geben. □

UPDATEHEAVY WORKLOAD Um die Fähigkeit der **PROTOTYPEN** zu messen große Mengen von Daten verändern zu können, wird ein **WORKLOAD** geladen der zu 95% aus **SCHREIBOPERATIONEN** besteht.

BIGTABLE weist während des **TESTS** das weit bessere **LATENZVERHALTEN** auf. Die Veränderung der **LATENZZEIT** entwickelt sich generell **LINEAR**. Wir wissen daß **RDBMS** Systeme **LESEOPTIMIERT** sind. Im Gegensatz dazu können **NoSQL** Systeme als **SCHREIBORIENTIERT** bezeichnet werden. Diese Systeme wurden darauf ausgelegt große Mengen von Daten schnell schreiben zu können.

Mehrere **ASPEKTE** von **NoSQL** Systemen zeichnen sich im Allgemeinen für das vorteilhafte **SCHREIBVERHALTEN** von **BIGTABLE** im Test aus:

- **DATENSHEMA: NoSQL DATENSCHMAS** sind darauf konzipiert, große Mengen von DATEN schnell schreiben zu können.

NoSQL DATENSCHMAS sind in der Regel **SCHEMALESS**. Müssen große Blöcke von KOMPLEXEN, ZUSAMMENHÄNGENDEN DATEN geschrieben werden, so erfolgt das im Normalfall in einer einzelnen **SCHREIBOPERATION**. Die DATEN werden dabei in einer einzelnen **ZELLE** oder in einem **DOKUMENT** abgelegt. Die Aufgabe des Programmieres ist es dabei die DATEN so zu strukturieren, daß Abhängigkeiten zwischen den DATEN hierarchisch modelliert werden können.

Um die **KONSISTENZ** der DATEN sicherzustellen, setzten **RDBMS** Systeme auf **TRANSAKTIONEN**. In **NoSQL SYSTEMEN** kennen wir das Konzept der Transaktion in der konservativen Interpretation nicht. In **NoSQL** Systemen gibt es kein **SCHEMA ON DEMAND**. Abhängigkeiten zwischen Daten werden auf der **ENTWURFSEBENE** und nicht auf der **SYSTEMEBENE** behandelt. □

- **KONSISTENZMODELL:** Das **BASE** Konsistenzmodell von **BIGTABLE** begünstigt das vorteilhafte **SCHREIBVERHALTEN** des Systems.

- **BASIC AVAILABILITY**
- **SOFT STATE**
- **EVENTUAL CONSISTENCY**

SOFT STATE als **KONSISTENZEIGENSCHAFT** bedeutet, daß die Daten im **SYSTEM** nicht zu jedem Zeitpunkt **KONSISTENT** sein müssen. **EVENTUAL CONSISTENCY** besagt daß die DATEN zu irgendeinem Zeitpunkt konsistent sein werden. Dieses **EPIDEMISCHE KONSISTENZMODELL** ermöglicht ein wesentlich effizienteres **SCHREIBVERHALTEN**. □

Analysieren wir das **SCHREIBVERHALTEN** des **MYSQL SYSTEMS** so bemerkt man daß nach einer gegebenen **LINEAREN** ENTWICKLUNG der **LATENZZEIT** bei steigender **SYSTEMLAST** die **FUNKTION** ein **EXPONENTIELLES** Wachstum erfährt. Das **SYSTEM** ist ab einem bestimmten Zeitpunkt nicht mehr in der Lage die **SYSTEMLAST** in einem vorgegebenen Rahmen zu verarbeiten.

Mehrer **ASPEKTE** von **RDBMS** Systemen sind für dieses Verhalten verantwortlich:

- **ACID KONSISTENZMODELL**
- **INDIZIERUNG DER DATEN**
- **RELATIONALE STRUKTUR DER DATEN**
- **MASTER/SLAVE SYSTEMSTRUKTUR**

SKALIERBARKEITS-TESTS

Als grundlegende NICHT FUNKTIONALE Anforderungen an unser SYSTEM haben wir ein **PERFORMANTES LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** des SYSTEMS gefordert. Ist ein SYSTEM **SKALIERBAR**, ist es in der Lage trotz wachsender **SYSTEM-LAST** den **SYSTEMSERVICE** adäquat aufrechtzuerhalten.

Durchführung der Tests:

- **WORKLOAD:** Zur Durchführung unsere TESTS wurde der **BALANCED WORKLOAD** geladen.
- **TEST:** Um die **SKALIERBARKEIT** des SYSTEMS messen zu können wird über eine bestimmte Zeit hinweg die Zahl der **SERVERKNOTEN** kontinuierlich erhöht. Gleichzeitig wird proportional dazu die **SERVERLAST** gesteigert.
- **MESSUNG:** Im Anschluß daran messen wir die entstehenden **LATENZZEITEN**.

MYSQL RESULTAT Die MESSUNGEN für den **MYSQL PROTOTYPEN** zeigen einen kleinen Anstieg in den **LATENZZEITEN**, der jedoch immer größer ausfällt und sich zunehmend **EXPONENTIELL** entwickelt.

RDBMS SYSTEME setzen im Gegensatz zu **NoSQL** Systemen auf den Betrieb mehrerer Servern durch **CLUSTERING**. Dieser Ansatz ist jedoch nur für eine beschränkte Zahl von Servern sinnvoll, da ab einem gewissen Punkt das System nicht mehr skaliert. Ein Grund für dieses Verhalten ist die Tatsache, daß **RDBMS SYSTEME** eine **MASTER/SLAVE ARCHITEKTUR** benötigen um die Konsistenz der Daten gewährleisten zu können. Die Notwendigkeit dieses Konsistenzverhaltens leitet sich wiederum aus den **ACID** Eigenschaften des Transaktionsmodells **RELATIONALER DATENBANKEN** her.

In diesem Umfeld erweist sich der **RELATIONALE TRANSAKTIONSANSATZ** mit dem Sperren von Datensätzen als Haupthindernis. Zum einen ist der **ZEIT- und KOMMUNIKATIONSAUFWAND** für eine Übereinkunft über das **SETZEN** und **AUFHEBEN** einer Sperre groß, zum anderen kann auf einen gesperrten Datensatz nicht lesend zugegriffen werden.

BIGTABLE RESULTAT Betrachten wir die **ENTWICKLUNG** der **MESSFUNKTION** für den **BIGTABLE PROTOTYPEN** so bemerken wir ein schwaches **LINEARES WACHSTUM** der Funktion. Der **PROTOTYP** zeigt in diesem Zusammenhang ein beinahe **IDEALES VERHALTEN**.

NoSQL als **TECHNOLOGIE** ist mit dem Anspruch angetreten **SKALIERENDE SYTEME** im **DATENVERARBEITUNGSBEREICH** zur Verfügung zu stellen. Der gesamte **TECHNOLOGIESTACK** wurde unter dem Augenmerk entwickelt **LINEAR** zu skalieren. Die grundlegende Idee des **NoSQL** Ansatzes ist es die **SOFTWAREKOMPONENTE** des Systems an die **HARDWAREKOMPONENTE** anzupassen. Dieser Ansatz ist es der **LINEARE SKALIERBARKEIT** ermöglicht.

TECHNOLOGIEKONZEPTE zur Umsetzung **LINEARER SKALIERBARKEIT** des PROTOTYPEN:

- **CONSISTENT HASHING**
- **BASE KONSISTENZMODELL**
- **MVCC**

ELASTIC SPEEDUP

Als **ELASTIC SPEEDUP** wird die EIGENSCHAFT eines **VERTEILTEN SYSTEMS** bezeichnet, zur Laufzeit **SERVER KNOTEN** in das SYSTEM zu integrieren ohne eine spürbare Erhöhung der **LATENZZEIT** zur LAUFZEIT zu bemerken. Ergibt sich eine erhöhte **LATENZZEIT** so ist das darauf zurückzuführen daß das System die neuen Komponenten erst adäquat initialisieren und integrieren muß. **SKALIERENDE SYSTEME** sollten so **ELASTISCH** wie möglich sein. □

BIGTABEL TEST Das SYSTEM verhält sich in der Regel **ELASTISCH**. Betrachten wir die **MESSUNGEN** für den PROTOTYPEN so merken wir ein gewisses **RAUSCHEN** in den Latenzzeiten, jedesmal wenn neue **SERVERKNOTEN** in das SYSTEM integriert werden. Das **RAUSCHEN** wird verursacht durch das **DATENMODELL** des Systems. Jedesmal wenn **SERVER KNOTEN** in das SYSTEM **INTEGRIERT** oder aus dem SYSTEM **GELÖSCHT** werden, beginnt der **CONSISTENT HASHING ALGORITHMUS** damit DATEN entsprechend umzukopieren und **REPLIKATE** zu erstellen. □

MYSQL TEST **MYSQL** wurde nicht darauf konzipiert **ELASTISCH** zu skalieren. □



Etablierung der Ergebnisse

FORSCHUNGSFRAGE

Wir sind mit der Frage angetreten ob es für eine bestimmte **KLASSE** von VERTEILTEN, DATENVERARBEITENDEN SYSTEMEN möglich ist ein **PERFORMANTES LAUFZEITVERHALTEN** und **SKALIERBARKEIT** zu fordern, bei der gleichzeitigen Verwendung eines **RDBMS SYSTEMS** als Kern dieses Systems.

KLASSE DER IN BERÜCKSICHTIGUNG ZU ZIEHENDEN SOFTWARESYSTEME Die für uns interessante **KLASSE VON SOFTWAREANWENDUNGEN** möchten wird durch die für sie CHARAKTERISTISCHEN EIGENSCHAFTEN beschreiben. Die gewünschte Klasse findet sich in der **DURCHSCHNITTMENGE** der durch folgende EIGENSCHAFTEN und ANFORDERUNGEN beschriebenen Softwaresysteme:

- **VERTEILTES SYSTEM**
- **BASIEREND AUF DEM HTTP PROTOKOLL**
- **UNTERSTÜTZT HTML 5**
- **HOCH SKALIERBAR**
- **DATENVERARBEITUNG IM TERRABYTE BEREICH**
- **PERFORMANTES LAUFZEITVERHALTEN** □

FORSCHUNGSFRAGE Die grundlegende Frage in diesem Zusammenhang ist damit:

SOLL ZUR ETABLIERUNG EINES VERTEILTEN WEBBASIERTEN CLIENT-SERVER SYSTEMS MIT MILLIONEN VON GLEICHZEITIGEN ANWENDERZUGRIFFEN UND DER NOTWENDIGKEIT DATEN IM PETA BYTE BEREICH ZU BEARBEITEN BEI GLEICHZEITIG PERFORMANTEM LAUFZEITVERHALTEN, EIN VERTIKAL ODER EIN HORIZONTAL SKALIERENDES SYSTEM EINGESETZT WERDEN? $\square$

METHODISCHES VORGEHEN

Um die **FORSCHUNGSFRAGE** beantworten zu können wurden die folgenden **ARBEITSSCHRITTE** durchlaufen:

- **TECHNOLOGIEVERGLEICH**
- **ENTWURF DER PROTOTYPEN**
- **AUSFÜHRUNG VON LASTENTESTS**
- **AUSWERTUNG**

TECHNOLOGIEVERGLEICH

CAP THEOREM Das **CAP THEOREM** etablierte für uns eine **KLASSIFIZIERUNG** von VERTEILTEN DATENVERARBEITENDEN SYSTEMEN. Aus der DISKUSSION des Theorems extrahierten wir 2 Klassen von Systemen die für einen Vergleich herangezogen werden sollten. (**CA SYSTEME** vs **CP SYSTEME**) □

CP- UND AP SYSTEME - NoSQL: **NoSQL** versteht sich als Sammlung diverser unterschiedlicher TECHNOLOGIEN und PARADIGMEN, die mit dem Anspruch angetreten sind, skalierbar und hoch performant zu sein. Eine **NoSQL** Datenbanksystem ist als DEZENTRALE STRUKTUR aufgebaut. Durch intelligente **HORIZONTALE SKALIERUNG** wird das Bild von nahezu uneingeschränkt verfügbarer Ressourcen geformt.

Die folgenden TECHNOLOGIEN werden in **NoSQL SYSTEMEN** eingesetzt:

- **KEY-VALUE DATENMODELL**
- **CONSISTENT HASHING**
- **MVCC**
- **SPALTENORIENTIERTES DATENMODELL**
- **DOKUMENTORIENTIERTES DATENMODELL**
- **MAP/REDUCE** □

CA SYSTEME - RDBMS: **RELATIONALE DATENBANKSYSTEME** wurden entwickelt um den wachsenden Ansprüchen einer neuen Generation von Applikationen gerecht zu werden. **RDBMS SYSTEME** wurden entwickelt um folgende Probleme im Datenverarbeitungsbereich zu lösen:

- **REDUNDANZ UND INKONSISTENZ** von Daten
- **BESCHRÄNKTE ZUGRIFFSMÖGLICHKEITEN**
- **HERAUSFORDERUNGEN DES MEHRBENUTZERSBETRIEBS**
- **VERLUST VON DATEN**
- **INTEGRITÄTSVERLETZUNG**
- **SICHERHEITSPROBLEME** □

NoSQL vs RDBMS SYSTEME RDBMS SYSTEME setzen im Gegensatz zu **NoSQL** Systemen auf den Betrieb mehrerer Servern durch **CLUSTERING**. Dieser Ansatz ist jedoch nur für eine beschränkte Zahl von SERVERN sinnvoll, da ab einem gewissen Punkt das System nicht mehr skaliert. Ein Grund für dieses Verhalten ist die Tatsache, daß **RDBMS SYSTEME** eine **MASTER/SLAVE ARCHITEKTUR** benötigen um die **KONSISTENZ** der Daten gewährleisten zu können. Die Notwendigkeit dieses Konsistenzverhaltens leitet sich wiederum aus den **ACID** Eigenschaften des Transaktionsmodells **RELATIONALER DATENBANKEN** her. Um die geforderten ANTWORTZEITEN erreichen zu können, mit mehreren Hunderttausend oder auch Millionen von Benutzern, bei gleichzeitiger **VERARBEITUNG** von DATENMENGEN im TERA- und PETA-BYTE-BEREICH, muß das **RELATIONALE KONSISTENZMODELL** hinterfragt werden. In diesem Umfeld erweist sich der **RELATIONALE TRANSAKTIONSANSATZ** mit dem Sperren von Datensätzen als Haupthindernis. Zum einen ist der Zeit- und Kommunikationsaufwand für eine Übereinkunft über das Setzen und Aufheben einer Sperre groß, zum anderen kann auf einen gesperrten Datensatz nicht lesend zugegriffen werden.

NoSQL versteht sich als Sammlung diverser unterschiedlicher TECHNOLOGIEN und PARADIGMEN, die entwickelt wurden, um die Restriktionen des **ACID** Konsistenzmodells zu umgehen, ohne auf **DATENKONSISTENZ** verzichten zu müssen.

Wir isolierten das **LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** als die entscheidenden NICHTFUNKTIONALE ANFORDERUNGEN in unserer Klasse von Softwareystemen. Ein performantes **LAUFZEITVERHALTEN** in Verbindung mit der Notwendigkeit, **MILLIONEN VON BENUTZERREQUESTS** zur selben Zeit zu bearbeiten, ist die Quadratur des Kreises, die es zu meistern gilt.

Eine Vielzahl von Requests gleichzeitig zu bearbeiten, birgt die Notwendigkeit, die Anwendung auf einen CLUSTER von Rechnern zu verteilen. Zur NICHTFUNKTIONALEN Anforderung des LAUFZEITVERHALTENS gesellt sich die Forderung nach einem SKALIERENDEN System. Das **LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** des Systems sind Größen, die in engem Zusammenhang mit anderen GRÖSSEN und KONZEPTEN stehen. □

RDBMS vs NoSQL: ÜBERSICHT

	CA SYSTEME	CP SYSTEME
SKALIERBARKEIT:	HORIZONTALE SKALIERBARKEIT	VERTIKALE SKALIERBARKEIT
KONSISTENZMODELL:	ACID	BASE
DATENSHEMA:	SCHEMALESS	SCHEMA
DATASTORE:	RDBMS	NoSQL

SKALIERBARKEIT **SKALIERBARKEIT** wird eingesetzt, um die **VERFÜGBARKEIT** von Anwendungen im industriellen Umfeld gewährleisten zu können. Ein **SOFTWARESYSTEM** wird als **SKALIERBAR** bezeichnet, wenn das **SYSTEM** in der Lage ist gleichzeitig zur wachsenden **SYSTEMLAST** den vereinbarten Dienst zu gewährleisten.

Generell wird unterschieden zwischen

- **HORIZONTALE/LINEARE SKALIERUNG**
- **VERTIKALER SKALIERUNG** □

KLASSIFIZIERUNG VON KONSISTENZMODELLEN

- **STRICT CONSISTENCY**
- **EVENTUAL CONSISTENCY**
- **FUNCTIONAL PARTITIONING** □

VERGLEICH ACID vs BASE

ACID	BASE
STRENGE KONSISTENZ	SCHWACHE KONSISTENZ
ISOLATION	VERFÜGBARKEIT
FOKUS AUF COMMIT	INETWA RICHTIGE ANTWORTEN
VERSCHACHTELTE TRANSAKTIONEN	SCHNELL
KONSERVATIV/PESSIMISTISCH	AGGRESSIV/OPTIMISTISCH
KOMPLIZIERTE EVOLUTION	EINFACHE EVOLUTION

ENTWURF DER PROTOTYPEN

Im Zuge der Arbeit haben wir einen Katalog von Anforderungen erarbeitet, der die Rahmenbedingungen zur Gestaltung der **ARCHITEKTUR** und technischen **UMSETZUNG** einer bestimmten Klasse von Applikationen definiert. Wir haben die unterschiedlichen **KONZEPTE** und **PARADIGMEN** kennengelernt, die für uns relevant sind, um das Problem zu verstehen und adäquat zu interpretieren.

Zur Etablierung unserer Ergebnisse werden 2 PROTOTYPEN erstellt und geprüft. Die PROTOTYPEN werden unterschiedlichen Paradigmen folgend entwickelt. Zum Vergleich werden die PROTOTYPEN identischen Lastentests unterzogen.

RDBMS vs NoSQL Beide Prototypen implementieren das **MVC MODELL**. Die Herausforderung in der Gestaltung beider Programme ist es, sowohl das **VIEW** als auch den **CONTROLLER** für beide identisch zu halten. Lediglich in der MODEL SCHICHT streben wir eine grundlegende Differenzierung im Sinne der Implementierung an.

Die PROTOTYPEN weisen eine grundlegende Differenzierung in ihrem **DATEN-**, **KONSISTENZ-** und **REPLIKATIONSMODELL** auf. Ein PROTOTYP wird den Prinzipien der **NoSQL** Ansätze folgend gestaltet, der andere nutzt die Theorie der verteilten **RELATIONALEN DATENBANKEN** für die DATENVERARBEITUNG.

AUSFÜHRUNG DER LASTENTESTS

LASTENTESTS werden eingesetzt, um das **LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** der Prototypen zu bewerten und zu verstehen. Das LAUFZEITVERHALTEN ist in diesem Sinne als die Summe der unterschiedlichen Paradigmen für **DATENKONSISTENZ**, **REPLIKATION** und des **DATENMODELLS** zu beurteilen.

Beide Prototypen werden denselben LASTENTESTS unterzogen. Der Vergleich der Ergebnisse der Lastentests gibt Auskunft über die wahre Beschaffenheit der unterschiedlichen Konzepte. Die gefundenen Ergebnisse helfen uns, die gewünschten Eigenschaften zu prüfen und finale Architekturentscheidung zu treffen. □

AUSWERTUNG

Wir haben ein **PERFORMANTES LAUFZEITVERHALTEN** und die **SKALIERBARKEIT** des SYSTEMS als primäre **NICHT FUNKTIONALE ANORDERUNG** an das SYSTEM gestellt. Die zentrale Frage ist ob es möglich ist ein solches System unter der Verwendung eines **RDBMS** Systems als Kern umzusetzen.

Betrachten wir die Ergebnisse für die **SKALIERBARKEIT** des MYSQL PROTOTYPEN bemerken wir daß das System nur **HORIZONTAL SKALIERT**. Unser Anspruch besteht jedoch darin ein SYSTEM zur Verfügung zu stellen, daß **LINEAR SKALIERT**. Diesen Anspruch erfüllt aber nur der **NoSQL PROTOTYP**. Wir müssen feststellen daß eine **RDBMS** Lösung für unser Problem nicht in Frage kommt. □

Bibliography

- [1] Tannenbaum, Steen, 2003. *Verteilte Systeme - Grundlagen und Paradigmen*, Pearson Studium.
- [2] Graham, Knuth, Patashnik, 1989. *Concrete Mathematics: A Foundation for Computer Science*, Addison-Wesley.
- [3] Heuer, 1997. *Objektdatenbanken - Konzepte, Sprachen, Architektur*, International Thomson Publishing.
- [4] Heuer, 2008. *Datenbanken - Konzepte und Sprachen*, mitp Verlag.
- [5] Saake, 2005. *Datenbanken: Implementierungstechniken*, mitp Verlag/Bonn.
- [6] R. A. Elmasri 2009. *Grundlagen von Datenbanksystemen*, Addison-Wesley.
- [7] Edlich, Friedland, Hampe 2010. *NoSQL- Einstieg in die Welt der nichtrelationalen Web 2.0 Anwendungen*, Hanser Verlag.
- [8] White, 2011. *Hadoop, The Definite Guide, 2ndEdition*, O'Reilly Media Inc.
- [9] Lam, 2010. *Hadoop in Action*, Manning Publications.
- [10] Date, 1994. *Mehrrechner-Datenbanksysteme*, Addison-Wesley.
- [11] Aviziens, Laprie, Randell, 2001. *Concepts of Dependability*, Research Report No 1145, LAAS-CNRS.
- [12] Brewer, 2000. *Towards Robust Distributed Systems*, Folien zur Keynote des 19 ACM SIGACT-SIGOPS Symposium
- [13] DeCandia, Hastorun, Jampani 2011. *Dynamo, Amazon's highly available Key-Value Store*, SIGOPS Operating Systems Rev. New York, pages 205-220
- [14] Karger, Lehman, Leighton 2011. *Consistent Hashing and Random Trees: Distributed Protocols for Relieving Hot Spots on the WWW*, STOC'97, Proceedings of the 29th anual ACM, Symposium on Theory of Computing, El Paso, Texas, USA, 1997, pages 654-663

- [15] Chang, Dean, Ghemawat, Burrows, Chandra 2006. *Bigtable: A Distributed Storage System for Structured Data*, Google inc., pages 205 - 218
- [16] Gilbert, Lynch, 2002. *Brewer's Conjecture and the Feasibility of Consistent, Available, Partition-Tolerant Web Services*, <http://lpd.epfl.ch/sgilbert/pubs/BrewersConjecture-SigAct.pdf>
- [17] <http://www.ocg.at/de/biber>
- [18] Wikipedia Beitrag, Pittimann, 2013. *HTTP*, http://de.wikipedia.org/wiki/Hypertext_Transfer_Protocol
- [19] Wikipedia Beitrag, Lippold, 2013. *HTML 5*, <http://de.wikipedia.org/wiki/HTML5>
- [20] SETI@home, <https://en.wikipedia.org/wiki/SETI@home>
- [21] Java Wiki <http://de.wikipedia.org/wiki/Java>
- [22] Avgeriou, Zdun. *Architectural Patterns Revisited – A Pattern Language*, <http://hydra.infosys.tuwien.ac.at/staff/zdun/publications/designSpacePatterns.pdf>
- [23] White Paper by DATASTAX CORPORATION 2013. <http://www.datastax.com/wp-content/uploads/2013/02/WP-Benchmarking-Top-NoSQL-Databases.pdf>

Index

- 11 g, 48
- ACID, 15
- Anforderungen, 5
 - nicht funktionale Anforderungen, 5
- Anforderungsanalyse, 25
 - Anwendungsanforderung, 26
 - HTTP Protokoll Implementierung, 26
 - Klasse von Softwareanwendungen, 25
 - Verteilte Systeme, 25
- AppEngine, 43
- Architecturalpattern
 - EAI/ESB, 66
 - ETL, 65
 - SOA, 66
- Architekturansatz
 - Patternsprache, 61
- BASE, 16
- Biber, 19
 - Aktorenliste, 20
 - Anforderungen, 23
 - Beschreibung, 19
 - Geschichte, 19
 - Usecases, 20
- Bigtable, 43
- CAP Theorem, 11
 - CA Systeme, 12
 - CP Systeme, 13
 - Klassifizierung, 12
- Cassandra, 38
- Column Family, 36
- Consistent Hashing, 31
- CouchDB, 40
- CQL, 39
- Database
 - AppEngine, 43
 - Cassandra, 38
 - CouchDB, 40
 - Hadoop-Hive, 41
 - MongoDB, 42
- Datenmodell
 - Consistent Hashing, 31
 - dokumentorientiert, 36
 - Key-Value, 35
 - Map/Reduce, 33
 - Schemafreiheit, 37
 - spaltenorientiert, 35
- Datenpersistenz
 - Functional Partitioning, 15
 - Konsistenzmodell, 15
- Datenverarbeitung, 6, 16
- Dokumentorientiertes Datenmodell, 36
- EAI/ESB, 66
- EER, 46
- Entity-Relationship Model, 46
- ETL, 65
- Eventual Consistency, 15
 - BASE, 16
- Fragestellung der Arbeit, 18
- Functional Partitioning, 15
- Hadoop, 7
- Hadoop-Hive, 41
- HQL, 41

- Key/Value Datenmodell, 35
- Konsistenz
 - Sale up, 6
- Konsistenzmodell, 6, 15
 - Eventual Consistency, 15
 - Strict Consistency, 15
- Kontext, 5
- Lastentest, 73
 - Framework, 74
 - Performance, 75
 - Skalierung, 75
 - Testkriterien, 75
 - Workload, 74
- Map/Reduce, 33
 - Mapper, 33
 - mapping phase, 33
 - Partitioning, 33
 - reduce phase, 34
 - Reducer, 34
 - Shuffling, 33
- Mapper, 33
- Methodisches Vorgehen, 8
 - Implementierung, 9
 - Lastentests, 10
 - Prototypen, 8
 - Resultat, 10
 - Softwareanwendungsklasse, 8
- MongoDB, 42
- Multiversion Concurrency Control, 32
- MVC Modell, 9
- MVCC, 32
- MySQL, 47
- nicht funktionale Anforderungen, 8
- NoSQL, 17, 38, 87
- NoSQL Datenbanken, 29
- Patternsprache, 61
- Problemstellung, 5
- Prototypen, 8
 - Architektur, 8
 - Implementierung, 9
- RDBMS, 45
- Reducer, 34
- Relationale Anfragesprache, 46
- Relationale Datenbanken, 45
 - Implementierungen, 47
 - Konsistenzmodell, 47
 - Prinzipien, 45
- Relationen, 46
- scale out, 29
- scale up, 29
- Schema, 46
- Schema On-Write, 6
- Schichtenmodell, 8
 - Abstraktion, 9
 - Modularisierung, 9
 - Schnittstellen, 9
- Home, 7
- Skalierbarkeit, 14
 - Horizontale Skalierung, 14
 - Vertikale Skalierung, 14
- SOA, 66
- Spaltenorientiertes Datenmodell, 35
- SQL, 46
- Strict Consistency, 15
 - ACID, 15
- Strukturierung
 - Kontext, 5
 - Methodisches Vorgehen, 8
 - Problemstellung, 5
- Technologievergleich, 11
 - CAP Theorem, 11
- Transaktionen, 47
- View
 - Adaption View, 62
 - Component Interaction View, 62
 - Data Centered View, 62
 - Distribution View, 62
 - Flow View, 62
 - Layered View, 62
- Web 2.0, 5, 6
- Wide Column Stores, 35