

1B Databases

Timothy G. Griffin

Computer Laboratory
University of Cambridge, UK

Lent 2017

Note on this transitional year

- For many years this has been a 1B course of 12 lectures held in Lent.
- In this academic year it has been transformed into a new Michaelmas course for 1A students (75-percent option) or (in the future) for 1B students (50-percent option).
- The new course is comprised of eight lectures and three practicals (“ticks”).
- Only in this transitional year is the course presented in both terms.
- Practical and ticks do not fit well into the current 1B pipeline!
- Lent’s 1B presentation will include repeats of the eight lectures from Michaelmas as well as 3 additional lectures covering the practical materials.
- One lecture (8 March, day of group project presentations) is canceled.

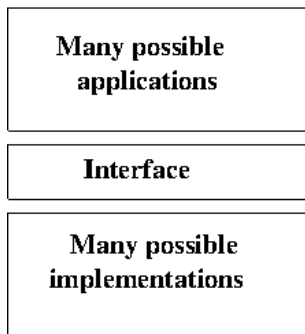
Outline

	date	topics
1	17/2	What is a Database Management System (DBMS)?
2	20/2	Entity-Relationship (ER) diagrams
3	22/2	Relational Databases ...
4	24/2	... and SQL
5	27/2	HyperSQL Practical Lecture
6	1/3	Some limitations of SQL ...
7	3/3	... that can be solved with Graph Database
8	6/3	Neo4j Practical Lecture
9	8/3	No Lecture!
10	17/2	Document-oriented Database
11	17/2	DoctorWho Practical Lecture
12	17/2	Cloud computing and distributed databases

Lecture 1

- What is a Database Management System (DBMS)?
- CRUD and ACID
- Three data models covered
 - ▶ Relational
 - ▶ Graph-oriented
 - ▶ Document-oriented
- Trade-offs in application design
- Trade-offs in DBMS design

Abstractions, interfaces, and implementations

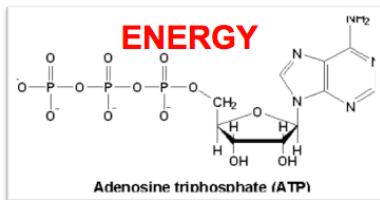


- An interface liberates application writers from low level details.
- An interface represents an abstraction of resources/services used by applications.
- In a perfect world, implementations can change without requiring changes to applications.
- Performance concerns often present a serious challenge to this idealised picture!

This is found everywhere, not just in computing ...



Evolution worked it out long long ago!



Photosynthesis



Cellular oxidation

What is a Database Management System (DBMS)?

Database
applications

Data model, query language,
programming API, ...

Query Engine, low-level
data representation, services

- This course will stress **data models** and **query languages**. We will not cover programming APIs or network APIs.
- A query engine knows about low-level details hidden by the interface(s). It uses this knowledge to optimize query evaluation.
- Primary service: persistent storage.
- Other services typically implemented:
 - ▶ CRUD operations,
 - ▶ ACID transactions.

CRUD operations

Create: Insert new data items into the database.

Read: Query the database.

Update: Modify objects in the database.

Delete: Remove data from the database.

ACID transactions

Atomicity: Either all actions of a transaction are carried out, or none are (even if the system crashes in the middle of a transaction).

Consistency: Every transaction applied to a consistent database leave it in a consistent state.

Isolation: Transactions are isolated, or protected, from the effects of other concurrently executed transactions.

Durability: If a transactions completes successfully, then its effects persist.

Implementing ACID transactions is one topic covered in Concurrent and Distributed Systems (1B).

This course looks at 3 data models

3 models

Relational Model: Data is stored in tables. SQL is the main query language.

Graph-oriented Model: Data is stored as a graph (nodes and edges). Query languages tend to have “path-oriented” capabilities.

Aggregate-oriented Model: Also called document-oriented database. Optimised for read-oriented databases.

The relational model has been the industry mainstay for the last 35 years. The other two models are representatives of an ongoing revolution in database systems often described under the (misleading) “NoSQL” banner.

This course uses 3 database systems



HyperSQL A Java-based relational DBMS. Query language is SQL.



Neo4j A Java-based graph-oriented DBMS. Query language is Cypher (named after a character in The Matrix).



DoctorWho A bespoke **document**-oriented DBMS. Stores JSON objects. Query language is Java.
(Under the hood and hidden from you: DoctorWho is implemented with Berkeley DB, a Key-value store.)

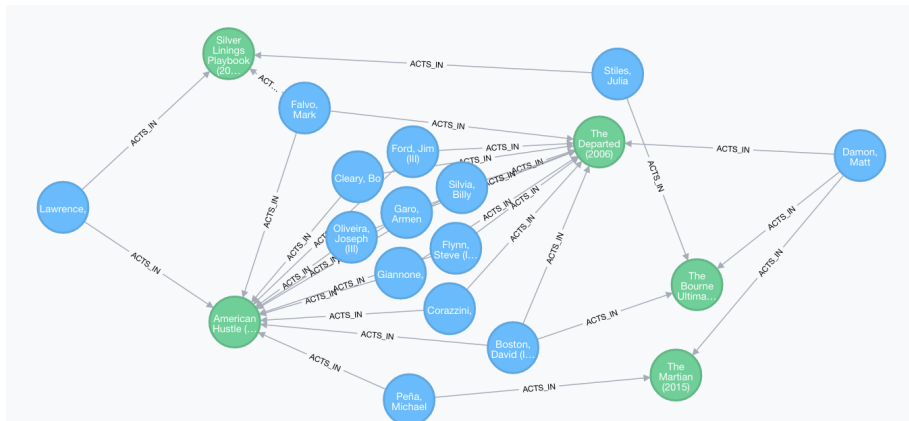
IMDb : Our data source



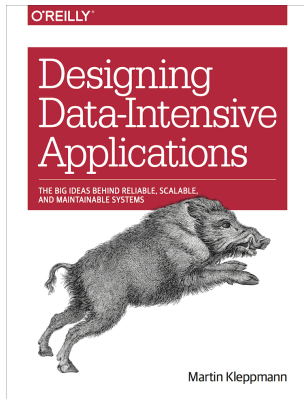
- We have processed raw data available from IMDb (plain text data files at <http://www.imdb.com/interfaces>)
- **Small instance:** 100 movies and 7771 people (actors, directors, etc)
- **Large instance:** 887,747 movies and 3,424,147 people
- For each of our three database systems we have generated a small and large database instance.
- Practical work should be done on the small instances. Large instances are for your entertainment and enlightenment.

Neo4j: Example of path-oriented query in Cypher (on small database)

```
match path=allshortestpaths(  
  (m:Person {name : 'Lawrence, Jennifer (III)'} )  
    [:ACTS_IN*]-  
  (n:Person {name : 'Damon, Matt'})  
)  
return path
```



Big thanks to Martin Kleppmann



- Martin is software engineer, entrepreneur, author and speaker.
- He co-founded Rapportive (acquired by LinkedIn in 2012) and Go Test It (acquired by Red Gate Software in 2009)
- Martin provided invaluable assistance with the preparation of this course and with the design and implementation of the infrastructure for the practicals
- <https://martin.kleppmann.com>

Why are there so many DBMS options?

- We will see that there is no one system that nicely solves all data problems.
- There are several **fundamental trade-offs** faced by application designers and system designers.
- A database engine might be optimised for a particular class of queries.

One important trade-off involves redundant data.

Redundant data

Informally, data in a database is **redundant** if it can be deleted and then reconstructed from the data remaining in the database.

A common trade-off: Query response vs. update throughput

Data redundancy is problematic for some applications

If a database supports many concurrent updates, then data redundancy leads to many problems, discussed in Lecture 4. If a database has little redundancy, then update throughput is typically better since transactions need only lock a few data items. This has been the traditional approach in the relational database world.

Data redundancy is highly desirable for some applications

In a low redundancy database, evaluation of complex queries can be very slow and require large amounts of computing power.

Precomputing answers to common queries (either fully or partially) can greatly speed up query response time. This introduces redundancy, but it may be appropriate for databases supporting applications that are read-intensive, with few or no data modifications. This is an approach common in aggregate-oriented databases.

Trade-offs often change as technology changes

Expect more dramatic changes in the coming decades ...



1956: A 5 megabyte hard drive



“768 Gig of RAM capacity”

Ideal for Virtualization + Database applications

Dual Xeon E5-2600 with 8 HD bays

CCSI, RSS004

A modern server

Lecture 2

- Conceptual modeling with Entity-Relationship (ER) diagrams
- Entities, attributes, and relationships
- Weak entities
- Cardinality of a relationship

Conceptual modeling with Entity-Relationship (ER) diagrams



Peter Chen

- It is very useful to have a **implementation independent** technique to describe the data that we store in a database.
- There are many formalisms for this, and we will use a popular one — Entity-Relationship (ER), due to Peter Chen (1976).
- The ER technique grew up around relational databases systems but it can help document and clarify design issues for any data model.

Entities capture things of interest



- **Entities** (squares) represent the nouns of our model
- **Attributes** (ovals) represent properties
- A **key** is an attribute whose value uniquely identifies an entity instance (here underlined)
- The **scope** of the model is limited — among the vast number of possible attributes that could be associated with a person, we are implicitly declaring that our model is concerned with only three.
- Very abstract, independent of implementation

Entity Sets (instances)

Instances of the Movie entity

- title: Gravity, year : 2013, id : 2945287
- title : The Bourne Ultimatum, year : 2007, id : 3527294

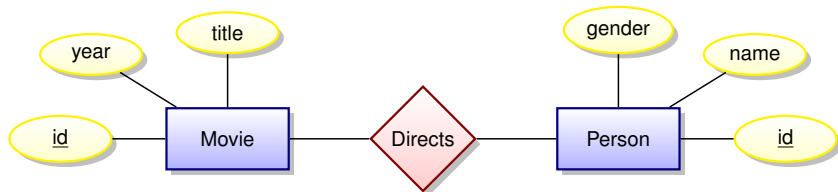
Instances of the Person entity

- name: Jennifer Lawrence, gender: female, id : 3018535
- name: Matt Damon, gender: male, id : 473946

Where do keys come from?

They are often automatically generated to be unique. Or they might be formed from some algorithm, like your CRSID. Q: Might some domains have natural keys (National Insurance ID)? A: Beware of using keys that are out of your control. The only safe thing to use as a key is something that is automatically generated in the database and only has meaning within that database.

Relationships



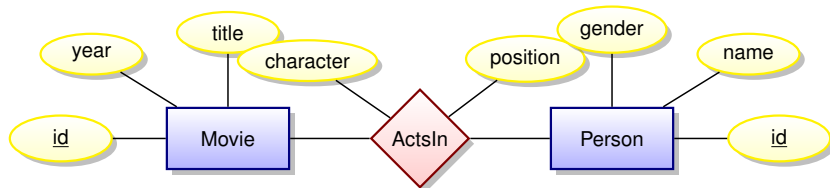
- Relationships (diamonds) represent the verbs of our domain.
- Relationships are between entities.

Relationship instances

Instances of the Directs relationship (ignoring entity attributes)

- Kathryn Bigelow directs The Hurt Locker
- Paul Greengrass directs The Bourne Ultimatum
- Steve McQueen directs 12 Years a Slave
- Karen Harley directs Waste Land
- Lucy Walker directs Waste Land
- João Jardim directs Waste Land

Relationships can have attributes



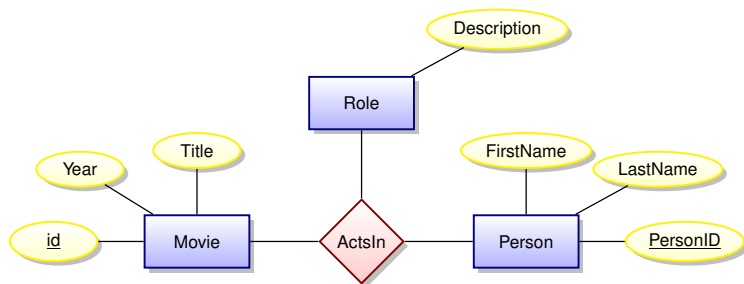
- Attribute **character** indicates the role played by a person
- Attribute **position** indicates the order listed in movie credits

Relationship instances

Instances of the ActsIn relationship (ignoring entity attributes)

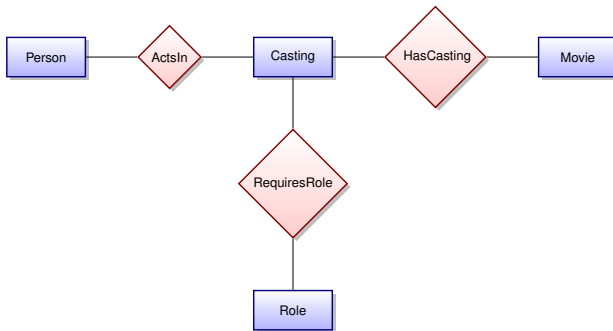
- Ben Affleck plays Tony Mendez in Argo, billing position 1
- Julie Deply plays Celine in Before Midnight, billing position 2
- Bradley Cooper plays Pat in Silver Linings Playbook, billing position 1
- Jennifer Lawrence plays Tiffany in Silver Linings Playbook, billing position 2
- Tim Allan plays Buzz Lightyear in Toy Story 3, billing position 2

Could **ActsIn** be modeled as a Ternary Relationship?



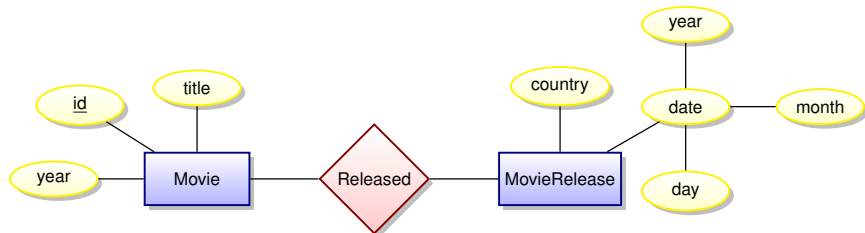
This might be a good model if our movies were in a restricted domain where roles have an independent existence. For example, suppose we are building a database of Shakespearean movies.

Can a ternary relationship be modeled with multiple binary relationships?



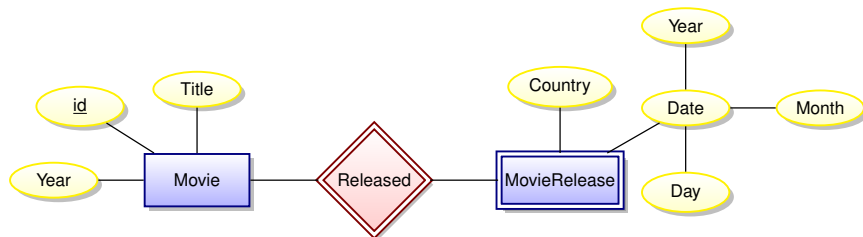
Is the **Casting** entity too artificial?

Attribute or entity with new relationship?



- Should release date be an attribute or an entity?
- The answer may depend on the **scope** of your data model.
- If all movies within your scope have at most one release date, then an attribute will work well.
- However, if your scope is global, then a movie can have different release dates in different countries.
- **Is there something strange about the MovieRelease?**

Weak entities



- **MovieRelease** is an example of a **weak entity**
- The existence of a weak entity depends on the existence of another entity. In this case, a release date exists only for a given movie.
- **Released** is called an **identifying relationship**

Cardinality of a relationship



The relation R is

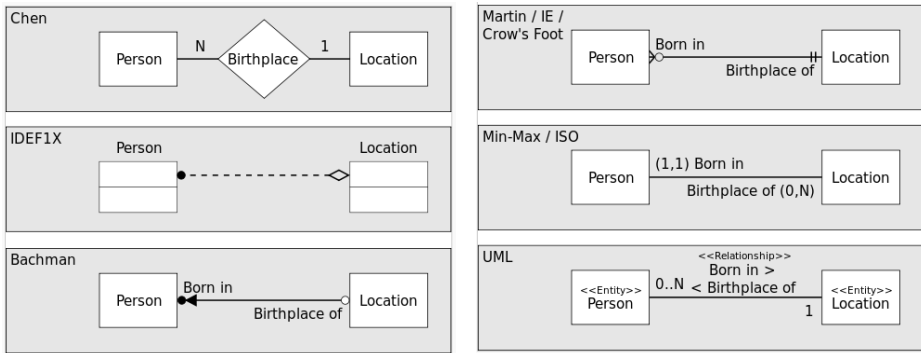
one-to-many: Every member of T is related to at most one member of S .

many-to-one: Every member of S is related to at most one member of T .

one-to-one: R is both many-to-one and one-to-many.

many-to-many: No constraint.

Diagrams can be annotated with cardinalities in many strange and wonderful ways ...



Various diagrammatic notations used to indicate a one-to-many relationship

https://en.wikipedia.org/wiki/Entity-relationship_model).

Note: We will not bother with these notations, but the concept of a relationship's cardinality is an important one.

Lectures 3 and 4

- The relational Model
- The Relational Algebra (RA)
- SQL
- Implementing an ER model in the relational model
- Update anomalies
- Avoiding redundancy

Still the dominant approach : Relational DBMSs

**your relational
application**

relational interface

**Database Management
System (DBMS)**

- In the 1970s you could not write a database application without knowing a great deal about the data's low-level representation.
- Codd's radical idea : give users a model of data and a language for manipulating that data which is completely independent of the details of its representation/implementation. That model is based on **mathematical relations**.
- This decouples development of the DBMS from the development of database applications.

Let's start with mathematical relations

Suppose that S and T are sets. The Cartesian product, $S \times T$, is the set

$$S \times T = \{(s, t) \mid s \in S, t \in T\}$$

A (binary) relation over $S \times T$ is any set R with

$$R \subseteq S \times T.$$

Database parlance

- S and T are referred to as **domains**.
- We are interested in **finite relations** R that can be stored!

n -ary relations

If we have n sets (domains),

$$S_1, S_2, \dots, S_n,$$

then an n -ary relation R is a set

$$R \subseteq S_1 \times S_2 \times \dots \times S_n = \{(s_1, s_2, \dots, s_n) \mid s_i \in S_i\}$$

Tabular presentation

1	2	...	n
x	y	...	w
u	v	...	s
$\vdots$	$\vdots$		$\vdots$
n	m	...	k

All data in a relational database is stored in **tables**. However, referring to columns by number can quickly become tedious!

Mathematical vs. database relations

Use named columns

- Associate a name, A_i (called an **attribute name**) with each domain S_i .
- Instead of tuples, use **records** — sets of pairs each associating an attribute name A_i with a value in domain S_i .

Column order does not matter

A database relation R is a **finite** set

$$R \subseteq \{ \{ (A_1, s_1), (A_2, s_2), \dots, (A_n, s_n) \} \mid s_i \in S_i \}$$

We specify R 's **schema** as $R(A_1 : S_1, A_2 : S_2, \dots, A_n : S_n)$.

Example

A relational schema

Students(**name**: string, **sid**: string, **age** : integer)

A relational instance of this schema

Students = {
 {(name, Fatima), (sid, fm21), (age, 20)},
 {(name, Eva), (sid, ev77), (age, 18)},
 {(name, James), (sid, jj25), (age, 19)}
}

Two equivalent tabular presentations

name	sid	age
Fatima	fm21	20
Eva	ev77	18
James	jj25	19

sid	name	age
fm21	Fatima	20
ev77	Eva	18
jj25	James	19

What is a (relational) database query language?

Input : a collection of
relation instances

Output : a single
relation instance

$$R_1, R_2, \dots, R_k \implies Q(R_1, R_2, \dots, R_k)$$

How can we express Q ?

In order to meet Codd's goals we want a query language that is high-level and independent of physical data representation.

There are **many** possibilities ...

The Relational Algebra (RA)

$Q ::=$	R	base relation
	$\sigma_p(Q)$	selection
	$\pi_{\mathbf{X}}(Q)$	projection
	$Q \times Q$	product
	$Q - Q$	difference
	$Q \cup Q$	union
	$Q \cap Q$	intersection
	$\rho_M(Q)$	renaming

- p is a simple boolean predicate over attributes values.
- $\mathbf{X} = \{A_1, A_2, \dots, A_k\}$ is a set of attributes.
- $M = \{A_1 \mapsto B_1, A_2 \mapsto B_2, \dots, A_k \mapsto B_k\}$ is a renaming map.
- A query Q must be **well-formed**: all column names of result are distinct. So in $Q_1 \times Q_2$, the two sub-queries cannot share any column names while in $Q_1 \cup Q_2$, the two sub-queries must share all column names.

SQL : a **vast** and **evolving** language

- Origins at IBM in early 1970's.
- SQL has grown and grown through many rounds of standardization :
 - ▶ ANSI: SQL-86
 - ▶ ANSI and ISO : SQL-89, SQL-92, SQL:1999, SQL:2003, SQL:2006, SQL:2008, SQL:2008
- SQL is made up of many sub-languages, including
 - ▶ Query Language
 - ▶ Data Definition Language
 - ▶ System Administration Language
- SQL will inevitably absorb many “NoSQL” features ...

Why talk about the Relational Algebra?

- Due to the RA's simple syntax and semantics, it can often help us better understand complex queries
- Tradition
- The RA lends itself to endlessly amusing tripos questions ...

Selection

R					$Q(R)$			
A	B	C	D	$\Rightarrow$	A	B	C	D
20	10	0	55		20	10	0	55
11	10	0	7		77	25	4	0
4	99	17	2					
77	25	4	0					

Q

RA $\sigma_{A>12}(R)$

SQL SELECT DISTINCT * FROM R WHERE R.A > 12

Projection

R					$Q(R)$	
A	B	C	D	$\Rightarrow$	B	C
20	10	0	55		10	0
11	10	0	7		99	17
4	99	17	2		25	4
77	25	4	0			

Q

RA $\pi_{B,C}(R)$

SQL SELECT DISTINCT B, C FROM R

Renaming

R					$Q(R)$			
A	B	C	D		A	E	C	F
20	10	0	55	$\Rightarrow$	20	10	0	55
11	10	0	7		11	10	0	7
4	99	17	2		4	99	17	2
77	25	4	0		77	25	4	0

Q

RA $\rho_{\{B \mapsto E, D \mapsto F\}}(R)$

SQL SELECT A, B AS E, C, D AS F FROM R

Union

R		S		$Q(R, S)$	
A	B	A	B	A	B
20	10	20	10	20	10
11	10	77	1000	11	10
4	99			4	99
				77	1000

Q

RA $R \cup S$

SQL (SELECT * FROM R) UNION (SELECT * FROM S)

Intersection

R		S		$Q(R)$	
A	B	A	B	A	B
20	10	20	10	20	10
11	10	77	1000		
4	99				

Q

RA $R \cap S$

SQL (SELECT * FROM R) INTERSECT (SELECT * FROM S)

Difference

R		S		$Q(R)$	
A	B	A	B	A	B
20	10	20	10	11	10
11	10	77	1000	4	99
4	99				

Q

RA $R - S$

SQL (SELECT * FROM R) EXCEPT (SELECT * FROM S)

Product

<i>R</i>		<i>S</i>		$Q(R, S)$			
<i>A</i>	<i>B</i>	<i>C</i>	<i>D</i>	<i>A</i>	<i>B</i>	<i>C</i>	<i>D</i>
20	10	14	99	20	10	14	99
11	10	77	100	20	10	77	100
4	99			11	10	14	99
				11	10	77	100
				4	99	14	99
				4	99	77	100

Q

RA $R \times S$

SQL SELECT A, B, C, D FROM R CROSS JOIN S

SQL SELECT A, B, C, D FROM R, S

Note that the RA product is not exactly the Cartesian product suggested by this notation!

Natural Join

First, a bit of notation

- We will often ignore domain types and write a relational schema as $R(\mathbf{A})$, where $\mathbf{A} = \{A_1, A_2, \dots, A_n\}$ is a set of attribute names.
- When we write $R(\mathbf{A}, \mathbf{B})$ we mean $R(\mathbf{A} \cup \mathbf{B})$ and implicitly assume that $\mathbf{A} \cap \mathbf{B} = \phi$.
- $u.[\mathbf{A}] = v.[\mathbf{A}]$ abbreviates $u.A_1 = v.A_1 \wedge \dots \wedge u.A_n = v.A_n$.

Natural Join

Given $R(\mathbf{A}, \mathbf{B})$ and $S(\mathbf{B}, \mathbf{C})$, we define the natural join, denoted $R \bowtie S$, as a relation over attributes $\mathbf{A}, \mathbf{B}, \mathbf{C}$ defined as

$$R \bowtie S \equiv \{t \mid \exists u \in R, v \in S, u.[\mathbf{B}] = v.[\mathbf{B}] \wedge t = u.[\mathbf{A}] \cup u.[\mathbf{B}] \cup v.[\mathbf{C}]\}$$

In the Relational Algebra:

$$R \bowtie S = \pi_{\mathbf{A}, \mathbf{B}, \mathbf{C}}(\sigma_{\mathbf{B}=\mathbf{B}'}(R \times \rho_{\vec{\mathbf{B}} \mapsto \vec{\mathbf{B}}'}(S)))$$

Join example

Students		
name	sid	cid
Fatima	fm21	cl
Eva	ev77	k
James	jj25	cl

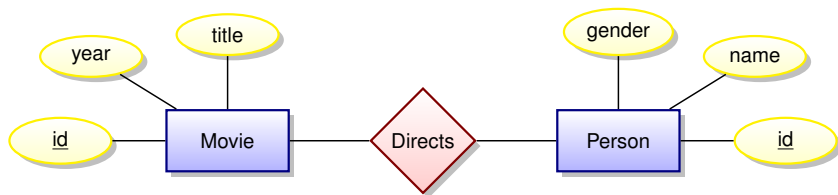
Colleges	
cid	cname
k	King's
cl	Clare
q	Queens'

⇒

Students ⋈ Colleges			
name	sid	cid	cname
Fatima	fm21	cl	Clare
Eva	ev77	k	King's
James	jj25	cl	Clare

(We will look at joins in SQL very soon ...)

How can we implement an ER model relationally?



- The ER model does not dictate implementation
- There are many options
- We will discuss some of the trade-offs involved

How about one table?

DirectsComplete

mid	title	year	pid	name	gender
2544956	12 Years a Slave (2013)	2013	1390487	McQueen, Steve (III)	male
2552840	4 luni, 3 saptamâni si 2 zile (2007)	2007	1486556	Mungiu, Cristian	male
2589088	Afghan Star (2009)	2009	3097222	Marking, Havana	female
2607939	American Hustle (2013)	2013	1835629	Russell, David O.	male
2611256	An Education (2009)	2009	3404232	Scherfig, Lone	female
2622261	Anvil: The Story of Anvil (2008)	2008	751015	Gervasi, Sacha	male
2626541	Argo (2012)	2012	16507	Affleck, Ben	male
2629853	Aruitemo aruitemo (2008)	2008	1133907	Koreeda, Hirokazu	male
...	...	...	...	...	...

This works, but

Anomalies caused by data redundancy

Insertion anomalies: How can we tell if a newly inserted record is consistent all other records records? We may want to insert a person without knowing if they are a director. We might want to insert a movie without knowing its director(s).

Deletion anomalies: We will wipe out information about people when last record is deleted from this table.

Update anomalies: What if an director's name is mis-spelled? We may update it correctly for one movie but not for another.

- A transaction implementing a conceptually simple update but containing checks to guarantee correctness may end up locking the entire table.
- Lesson: In a database supporting many concurrent updates we see that data redundancy can lead to complex transactions and low write throughput.

A better idea : break tables down in order to reduce redundancy

Movies

id	title	year
2544956	12 Years a Slave (2013)	2013
2552840	4 luni, 3 saptamâni si 2 zile (2007)	2007
2589088	Afghan Star (2009)	2009
2607939	American Hustle (2013)	2013
2611256	An Education (2009)	2009
2622261	Anvil: The Story of Anvil (2008)	2008
2626541	Argo (2012)	2012
2629853	Aruitemo aruitemo (2008)	2008
...	...	...

People

id	name	gender
1390487	McQueen, Steve (III)	male
1486556	Mungiu, Cristian	male
3097222	Marking, Havana	female
1835629	Russell, David O.	male
3404232	Scherfig, Lone	female
751015	Gervasi, Sacha	male
16507	Affleck, Ben	male
1133907	Koreeda, Hirokazu	male
...	...	...

What about the relationship?

Directs

movie_id	person_id
2544956	1390487
2552840	1486556
2589088	3097222
2607939	1835629
2611256	3404232
2622261	751015
2626541	16507
2629853	1133907
...	...

(No, this relation does not actually exist in our IMDb databases — more on that later ...)

We can recover DirectsComplete using a relational query

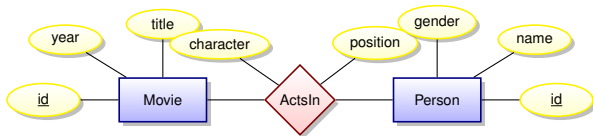
The SQL query

```
select movies.id as mid, title, year,  
       people.id as pid, name, gender  
from movies  
join directs on movie_id = movies.id  
join people on people.id = person_id
```

might return something like

MID	TITLE	YEAR	PID	NAME	GENDER
2544956	12 Years a Slave (2013)	2013	1390487	McQueen, Steve (III)	male
2552840	4 luni, 3 saptamâni si 2 zile (2007)	2007	1486556	Mungiu, Cristian	male
2589088	Afghan Star (2009)	2009	3097222	Marking, Havana	female
2607939	American Hustle (2013)	2013	1835629	Russell, David O.	male
2611256	An Education (2009)	2009	3404232	Scherfig, Lone	female
2622261	Anvil: The Story of Anvil (2008)	2008	751015	Gervasi, Sacha	male
2626541	Argo (2012)	2012	16507	Affleck, Ben	male
2629853	Aruitemo aruitemo (2008)	2008	1133907	Koreeda, Hirokazu	male
...	...	...	...	...	...

In a similar way with the ActsIn relationship



The ActsIn Relationship requires attributes

movie_id	person_id	character	position
2544956	146271	Judge Turner	4
2544956	2460265	Mistress Ford	32
2544956	173652	Mr. Moon	9
2544956	477824	Tibeats	35
2544956	256114	Edward	42
2544956	2826281	Tea Seller	NULL
...	...	...	...

More on NULL soon ...

We can recover all information for the ActsIn relation

The SQL query

```
select movies.id as mid, title, year,  
       people.id as pid, name, character, position  
from movies  
join actsin on movie_id = movies.id  
join people on people.id = person_id
```

might return something like

MID	TITLE	YEAR	PID	NAME	CHARACTER	POSITION
-----	-----	-----	-----	-----	-----	-----
2544956	12 Years a Slave (2013)	2013	146271	Batt, Bryan	Judge Turner	4
2544956	12 Years a Slave (2013)	2013	2460265	Bennett, Liza J.	Mistress Ford	32
2544956	12 Years a Slave (2013)	2013	173652	Bentley, Tony (I)	Mr. Moon	9
2544956	12 Years a Slave (2013)	2013	477824	Dano, Paul	Tibeats	35
2544956	12 Years a Slave (2013)	2013	256114	Bright, Gregory	Edward	42
2544956	12 Years a Slave (2013)	2013	2826281	Haley, Emily D.	Tea Seller	NULL
...	...	...	...	...	...	...

Observations

- Both ER entities and ER relationships are implemented as tables.
- We call them tables rather than relations to avoid confusion!
- Good: We avoid many update anomalies by breaking tables into smaller tables.
- Bad: We have to work hard to combine information in tables (joins) to produce interesting results.

What about consistency/integrity of our relational implementation?

How can we ensure that the table representing an ER relation really implements a relationship? Answer : we use **keys** and **foreign keys**.

Key Concepts

Relational Key

Suppose $R(\mathbf{X})$ is a relational schema with $\mathbf{Z} \subseteq \mathbf{X}$. If for any records u and v in any instance of R we have

$$u.[\mathbf{Z}] = v.[\mathbf{Z}] \implies u.[\mathbf{X}] = v.[\mathbf{X}],$$

then $\mathbf{Z}$ is a **superkey for R** . If no proper subset of $\mathbf{Z}$ is a superkey, then $\mathbf{Z}$ is a **key for R** . We write $R(\underline{\mathbf{Z}}, \mathbf{Y})$ to indicate that $\mathbf{Z}$ is a key for $R(\mathbf{Z} \cup \mathbf{Y})$.

Note that this is a **semantic** assertion, and that a relation can have multiple keys.

Foreign Keys and Referential Integrity

Foreign Key

Suppose we have $R(\underline{\mathbf{Z}}, \mathbf{Y})$. Furthermore, let $S(\mathbf{W})$ be a relational schema with $\mathbf{Z} \subseteq \mathbf{W}$. We say that $\mathbf{Z}$ represents a **Foreign Key in S for R** if for any instance we have $\pi_{\mathbf{Z}}(S) \subseteq \pi_{\mathbf{Z}}(R)$. Think of these as (logical) pointers!

Referential integrity

A database is said to have **referential integrity** when all foreign key constraints are satisfied.

A relational representation

A relational schema

$ActsIn(\underline{movie_id}, \underline{person_id})$

With **referential integrity constraints**

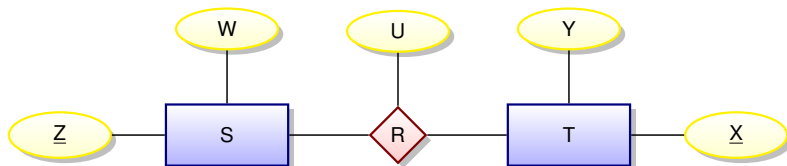
$$\pi_{movie_id}(ActsIn) \subseteq \pi_{id}(Movies)$$

$$\pi_{person_id}(ActsIn) \subseteq \pi_{id}(People)$$

Foreign Keys in SQL

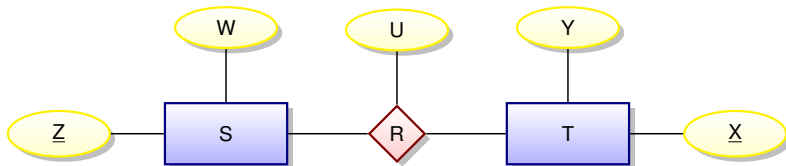
```
create table ActsIn (  
    movie_id int not NULL,  
    person_id int not NULL,  
    character varchar(255),  
    position integer,  
  
    primary key (movie_id, person_id),  
    constraint actsin_movie  
        foreign key (movie_id)  
        references Movies(id),  
    constraint actsin_person  
        foreign key (person_id)  
        references People(id))
```

Relationships to Tables (the “clean” approach)



Relation R is	Schema
many to many ($M : N$)	$R(\underline{X}, \underline{Z}, U)$
one to many ($1 : M$)	$R(\underline{X}, Z, U)$
many to one ($M : 1$)	$R(X, \underline{Z}, U)$
one to one ($1 : 1$)	$R(\underline{X}, Z, U)$ and/or $R(X, \underline{Z}, U)$

Implementation can differ from the “clean” approach

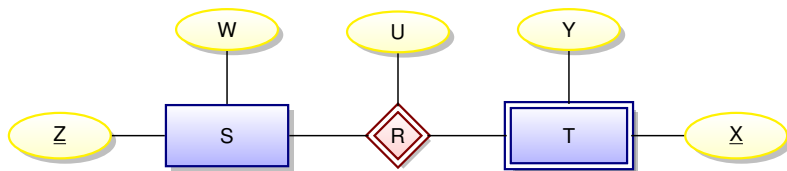


Suppose R is one to many

Rather than implementing a new table $R(\underline{X}, Z, U)$ we could expand table $T(\underline{X}, Y)$ to $T(\underline{X}, Y, Z, U)$ and allow the Z and U columns to be NULL for those rows in T not participating in the relationship.

Pros and cons?

Implementing weak entities

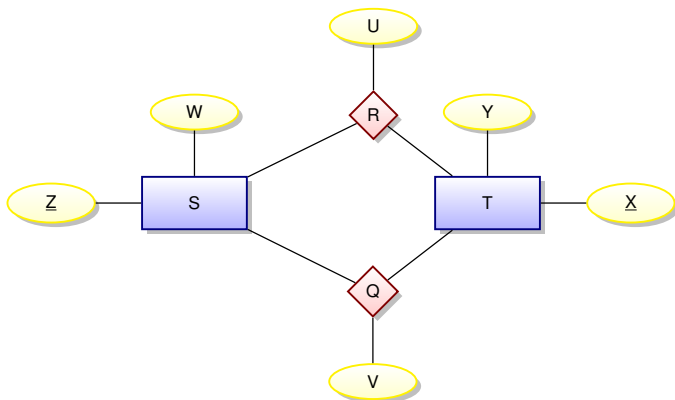


This is always a one to many relationship!

- Notice that **all** rows of T must participate in the relationship.
- The expanded $T(\underline{X}, Y, Z, U)$ is even more compelling.
- We might drop the keys $\underline{X}$ from T resulting in $T(\underline{Y}, \underline{Z}, \underline{U})$.
- This is exactly what we have done in our movies database with these tables: release_dates, certificates, color_info, genres, keywords, languages, locations, running_times.

Implementing multiple relationships into a single table?

Suppose we have two many-to-many relationships:



Implementing multiple relationships into a single table?

Rather than using two tables

$$\begin{array}{l} R(\underline{X}, \underline{Z}, U) \\ Q(\underline{X}, \underline{Z}, V) \end{array}$$

we might squash them into a single table

$$RQ(\underline{X}, \underline{Z}, \textit{type}, U, V)$$

using a tag $\textit{domain}(\textit{type}) = \{\mathbf{r}, \mathbf{q}\}$ (for some constant values r and q).

- represent an R -record (x, z, u) as an RQ -record $(x, z, \mathbf{r}, u, \text{NULL})$
- represent an Q -record (x, z, v) as an RQ -record $(x, z, \mathbf{q}, \text{NULL}, v)$

Redundancy alert!

If we now the value of the $\textit{type}$ column, we can compute the value of either the U column or the V column!

We have stuffed 9 relationships into the `credits` table!

```
SELECT type, count(*) AS total
FROM credits
GROUP BY type
ORDER BY total DESC;
```

This query produces the output (on the small database):

TYPE	TOTAL
-----	-----
actor	6913
producer	905
writer	212
editor	134
director	113
cinematographer	108
composer	101
production_designer	83
costume_designer	72

Was this a good idea?

Discuss!

CONTEST!!!

Explore the (small or large) relational database using SQL and discover surprising/amusing/strange/wonderful things in the data!

Send findings (queries and results) to `tgg22@cam.ac.uk`.

PRIZES!!!



Lecture 5 : HyperSQL Practical Lecture

Before lecture, please attempt **Getting started with HyperSQL** and **HyperSQL exercises** found on the course web page.

Lectures 6 and 7



- What is a database index?
- Two complications for SQL semantics
 - ▶ Multi-sets (bags)
 - ▶ NULL values
- **Kevin Bacon!**
- Transitive closure of a relation
- Problems computing a transitive closure in relational databases
- Graph-oriented databases: optimised for computing transitive closure
- Neo4j
- Bacon number with Neo4j

Complexity of a JOIN?

Given tables $R(\mathbf{A}, \mathbf{B})$ and $S(\mathbf{B}, \mathbf{C})$, how much work is required to compute the join $R \bowtie S$?

```
// Brute force appaoch:
// scan R
for each (a, b) in R {
    // scan S
    for each (b', c) in S {
        if b = b' then create (a, b, c) ...
    }
}
```

Worst case: requires on the order of $|R| \times |S|$ steps. But note that on each iteration over R , there may be only a very small number of matching records in S — only one if R has a foreign key into S .

What is a database index?

An **index** is a data structure — created and maintained within a database system — that can greatly reduce the time needed to locate records.

```
// scan R
for each (a, b) in R {
    // don't scan S, use an index
    for each s in S-INDEX-ON-B(b) {
        create (a, b, s.c) ...
    }
}
```

In 1A Algorithms you will see a few of the data structures used to implement database indices (search trees, hash tables, and so on).

Remarks

Typical SQL commands for creating and deleting an index:

```
CREATE INDEX index_name on S (B)
```

```
DROP INDEX index_name
```

- There are many types of database indices and the commands for creating them can be complex.
- Index creation is not defined in the SQL standards.
- **While an index can speed up reads, it will slow down updates. This is one more illustration of a fundamental database tradeoff.**
- The tuning of database performance using indices is a fine art.
- In some cases it is better to store read-oriented data in a separate database optimised for that purpose.

Why the `distinct` in the SQL?

The SQL query

```
select B, C from R
```

will produce a bag (multiset)!

R					$Q(R)$		
A	B	C	D	$\Rightarrow$	B	C	
20	10	0	55		10	0	***
11	10	0	7		10	0	***
4	99	17	2		99	17	
77	25	4	0		25	4	

SQL is actually based on multisets, not sets.

Why Multisets?

Duplicates are important for aggregate functions (min, max, ave, count, and so on). These are typically used with the **GROUP BY** construct.

sid	course	mark
ev77	databases	92
ev77	spelling	99
tgg22	spelling	3
tgg22	databases	100
fm21	databases	92
fm21	spelling	100
jj25	databases	88
jj25	spelling	92

group $\Rightarrow$ by

course	mark
spelling	99
spelling	3
spelling	100
spelling	92

course	mark
databases	92
databases	100
databases	92
databases	88

Visualizing the aggregate function **min**

course	mark
spelling	99
spelling	3
spelling	100
spelling	92

course	mark
databases	92
databases	100
databases	92
databases	88

min(**mark**)
 $\Rightarrow$

course	min(mark)
spelling	3
databases	88

In SQL

```
select course,  
        min(mark) ,  
        max(mark) ,  
        avg(mark)  
from marks  
group by course;
```

course	min(mark)	max(mark)	avg(mark)
databases	88	100	93.0000
spelling	3	100	73.5000

What is NULL?

- NULL is a **place-holder**, not a value!
- NULL is not a member of any domain (type),
- This means we need three-valued logic.

Let $\perp$ represent **we don't know!**

$\wedge$	T	F	$\perp$
T	T	F	$\perp$
F	F	F	F
$\perp$	$\perp$	F	$\perp$

$\vee$	T	F	$\perp$
T	T	T	T
F	T	F	$\perp$
$\perp$	T	$\perp$	$\perp$

$\neg$	$\neg V$
T	F
F	T
$\perp$	$\perp$

NULL can lead to unexpected results

```
select * from students;
```

sid	name	age
ev77	Eva	18
fm21	Fatima	20
jj25	James	19
ks87	Kim	NULL

```
select * from students where age <> 19;
```

sid	name	age
ev77	Eva	18
fm21	Fatima	20

The ambiguity of NULL

Possible interpretations of NULL

- There is a value, but we don't know what it is.
- No value is applicable.
- The value is known, but you are not allowed to see it.
- ...

A great deal of semantic muddle is created by conflating all of these interpretations into one non-value.

On the other hand, introducing distinct NULLs for each possible interpretation leads to very complex logics ...

SQL's NULL has generated endless controversy

C. J. Date [D2004], Chapter 19

“Before we go any further, we should make it very clear that in our opinion (and in that of many other writers too, we hasten to add), NULLs and 3VL are and always were a serious mistake and have no place in the relational model.”

In defense of Nulls, by Fesperman

“[...] nulls have an important role in relational databases. To remove them from the currently **flawed** SQL implementations would be throwing out the baby with the bath water. On the other hand, the **flaws** in SQL should be repaired immediately.” (See <http://www.firstsql.com/idefend.htm>.)

Flaws? One example of SQL's inconsistency

With our small database, the query

```
SELECT note FROM credits WHERE note IS NULL;
```

returns 4232 records of NULL.

The expression `note IS NULL` is either true or false — true when `note` is the NULL value, false otherwise.

Flaws? One example of SQL's inconsistency (cont.)

Furthermore, the query

```
SELECT note, count(*) AS total
FROM credits
WHERE note IS NULL GROUP BY note;
```

returns a single record

```
note total
---- -
NULL 4232
```

This seems to mean that `NULL` is equal to `NULL`. But recall that `NULL = NULL` returns `NULL`!

Bacon Number

- Kevin Bacon has Bacon number 0.
- Anyone acting in a movie with Kevin Bacon has Bacon number 1.
- For any other actor, their bacon number is calculated as follows. Look at all of the movies the actor acts in. Among all of the associated co-actors, find the smallest Bacon number k . Then the actor has Bacon number $k + 1$.

Let's try to calculate Bacon numbers using SQL!

Mathematical relations, again

Given two binary relations

$$\begin{aligned} R &\subseteq S \times T \\ Q &\subseteq T \times U \end{aligned}$$

we can define their **composition** $Q \circ R \subseteq S \times U$ as

$$Q \circ R \equiv \{(s, u) \mid \exists t \in T, (s, t) \in R \wedge (t, u) \in Q\}$$

Partial functions as relations

- A (partial) function $f \in S \rightarrow T$ can be thought of as a binary relations where $(s, t) \in f$ if and only if $t = f(s)$.
- Suppose R is a relation where if $(s, t_1) \in R$ and $(s, t_2) \in R$, then it follows that $t_1 = t_2$. In this case R represents a (partial) function.
- Given (partial) functions $f \in S \rightarrow T$ and $g \in T \rightarrow U$ their **composition** $g \circ f \in S \rightarrow U$ is defined by $(g \circ f)(s) = g(f(s))$.
- Note that the definition of $\circ$ for relations and functions is equivalent for relations representing functions.

Since we could write $Q \circ R$ as $R \bowtie_{2=1} Q$ we can see that **joins are a generalisation of function composition!**

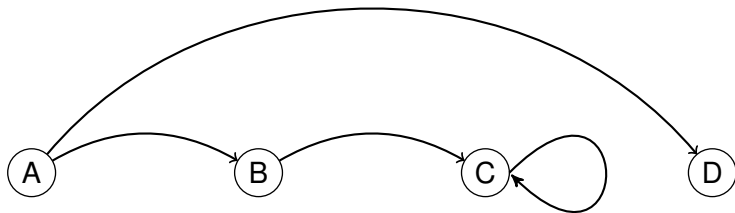
Directed Graphs

- $G = (V, A)$ is a **directed graph**, where
- V a finite set of **vertices** (also called **nodes**).
- A is a binary relation over V . That is $A \subseteq V \times V$.
- If $(u, v) \in A$, then we have an **arc** from u to v .
- The arc $(u, v) \in A$ is also called a directed edge, or a **relationship of u to v** .

Drawing directed graphs

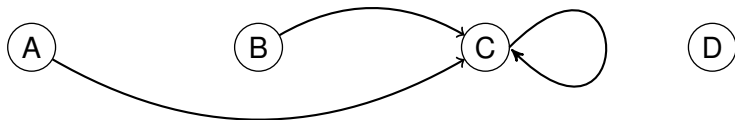
A directed graph

- $V = \{A, B, C, D\}$
- $A = \{(A, B), (A, D), (B, C), (C, C)\}$



Composition example

$$A \circ A = \{(A, C), (B, C), (C, C)\}$$



Elements of $A \circ A$ represent paths of length 2

- $(A, C) \in A \circ A$ by the path $A \rightarrow B \rightarrow C$
- $(B, C) \in A \circ A$ by the path $B \rightarrow C \rightarrow C$
- $(C, C) \in A \circ A$ by the path $C \rightarrow C \rightarrow C$

Iterated composition, and paths

Suppose R is a binary relation over S , $R \subseteq S \times S$. Define **iterated composition** as

$$\begin{aligned} R^1 &\equiv R \\ R^{n+1} &\equiv R \circ R^n \end{aligned}$$

Let $G = (V, A)$ be a directed graph. Suppose $v_1, v_2, \dots, v_{k+1}$ is a sequence of vertices. Then this sequence represents a **path in G of length k** when $(v_i, v_{i+1}) \in A$, for $i \in \{1, 2, \dots, k\}$. We will often write this as

$$v_1 \rightarrow v_2 \rightarrow \dots v_k$$

Observation

If $G = (V, A)$ is a directed graph, and $(u, v) \in A^k$, then there is at least one path in G from u to v of length k . Such paths may contain loops.

Shortest path

Definition of R -distance (hop count)

Suppose $s_0 \in \pi_1(R)$ (that is there is a pair $(s_0, s_1) \in R$).

- The distance from s_0 to s_0 is 0.
- If $(s_0, s_1) \in R$, then the distance from s_0 to s_1 is 1.
- For any other $s' \in \pi_2(R)$, the distance from s_0 to s' is the least n such that $(s_0, s') \in R^n$.

We will think of the Bacon number as an R -distance where s_0 is Kevin Bacon. But what is R ?

Let R be the co-actor relation

```
SELECT distinct p1.name as name1,  
                p1.id as id1,  
                p2.name as name2,  
                p2.id as id2  
FROM people AS p1  
JOIN credits AS c1 ON c1.person_id = p1.id  
JOIN credits AS c2 ON c2.movie_id = c1.movie_id  
JOIN people AS p2 ON p2.id = c2.person_id  
where c1.type = 'actor'  
      AND c2.type = 'actor';
```

On our small database the co-actor relation contains 748,420 rows.

SQL : Bacon number 1

```
SELECT distinct name
FROM credits AS c1
JOIN credits AS c2 ON c2.movie_id = c1.movie_id
JOIN people      ON id = c2.person_id
WHERE c1.person_id = 107303 AND id <> 107303
      AND c1.type = 'actor' AND c2.type = 'actor';
```

Compute names of all actors with Bacon number 1. Small database: 101 rows. (Kevin Bacon's id is 107303.)

SQL : Bacon number 2

```
SELECT distinct name
FROM credits AS c1
JOIN credits AS c2 ON c2.movie_id = c1.movie_id
JOIN credits AS c3 ON c3.person_id = c2.person_id
JOIN credits AS c4 ON c4.movie_id = c3.movie_id
JOIN people      ON id = c4.person_id
WHERE c1.person_id = 107303 AND id <> 107303
      AND c1.type = 'actor' AND c2.type = 'actor'
      AND c3.type = 'actor' AND c4.type = 'actor'
EXCEPT
SELECT distinct name
FROM credits AS c1
JOIN credits AS c2 ON c2.movie_id = c1.movie_id
JOIN people      ON id = c2.person_id
WHERE c1.person_id = 107303 AND id <> 107303
      AND c1.type = 'actor' AND c2.type = 'actor';
```

Compute names of all actors with Bacon number 2. Small database:
1387 rows.

SQL : Bacon number 3

```
SELECT distinct name
FROM credits AS c1
JOIN credits AS c2 ON c2.movie_id = c1.movie_id
JOIN credits AS c3 ON c3.person_id = c2.person_id
JOIN credits AS c4 ON c4.movie_id = c3.movie_id
JOIN credits AS c5 ON c5.person_id = c4.person_id
JOIN credits AS c6 ON c6.movie_id = c5.movie_id
JOIN people      ON id = c6.person_id
WHERE c1.person_id = 107303 AND id <> 107303
      AND c1.type = 'actor' AND c2.type = 'actor'
      AND c3.type = 'actor' AND c4.type = 'actor'
      AND c5.type = 'actor' AND c6.type = 'actor'
EXCEPT
(
  (SELECT distinct name
   FROM credits AS c1
   JOIN credits AS c2 ON c2.movie_id = c1.movie_id
   JOIN credits AS c3 ON c3.person_id = c2.person_id
   JOIN credits AS c4 ON c4.movie_id = c3.movie_id
   JOIN people      ON id = c4.person_id
   WHERE c1.person_id = 107303 AND id <> 107303
        AND c1.type = 'actor' AND c2.type = 'actor'
        AND c3.type = 'actor' AND c4.type = 'actor'
   EXCEPT
   SELECT distinct name
   FROM credits AS c1
   JOIN credits AS c2 ON c2.movie_id = c1.movie_id
   JOIN people      ON id = c2.person_id
   WHERE c1.person_id = 107303 AND id <> 107303
        AND c1.type = 'actor' AND c2.type = 'actor')
 UNION
  (SELECT distinct name
   FROM credits AS c1
   JOIN credits AS c2 ON c2.movie_id = c1.movie_id
   JOIN people AS p ON id = c2.person_id
   WHERE c1.person_id = 107303 AND id <> 107303
        AND c1.type = 'actor' AND c2.type = 'actor'));
```

Compute names of all actors
with Bacon number 3. Small
database: 3186 rows.

SQL : Bacon number 3, simplified

```
SELECT distinct name
FROM credits AS c1
JOIN credits AS c2 ON c2.movie_id = c1.movie_id
JOIN credits AS c3 ON c3.person_id = c2.person_id
JOIN credits AS c4 ON c4.movie_id = c3.movie_id
JOIN credits AS c5 ON c5.person_id = c4.person_id
JOIN credits AS c6 ON c6.movie_id = c5.movie_id
JOIN people      ON id = c6.person_id
WHERE c1.person_id = 107303 AND id <> 107303
      AND c1.type = 'actor' AND c2.type = 'actor'
      AND c3.type = 'actor' AND c4.type = 'actor'
      AND c5.type = 'actor' AND c6.type = 'actor'
EXCEPT []
SELECT distinct name
FROM credits AS c1
JOIN credits AS c2 ON c2.movie_id = c1.movie_id
JOIN credits AS c3 ON c3.person_id = c2.person_id
JOIN credits AS c4 ON c4.movie_id = c3.movie_id
JOIN people      ON id = c4.person_id
WHERE c1.person_id = 107303 AND id <> 107303
      AND c1.type = 'actor' AND c2.type = 'actor'
      AND c3.type = 'actor' AND c4.type = 'actor';
```

That last query was of the form

A EXCEPT ((B EXCEPT C) UNION C)

so we can simplify to

A EXCEPT B

A very unscientific experiment

Time (in seconds) to compute Bacon number $n \in \{2, 3, 4\}$ on my shabby laptop (includes database start-up time).

	2	3	4
HyperSQL	3	9	OOM

OOM = Out Of Memory.

This probably says more about HyperSQL than it does about Relational databases in general!

Transitive closure

Suppose R is a binary relation over S , $R \subseteq S \times S$. The **transitive closure of R** , denoted R^+ , is the smallest binary relation on S such that $R \subseteq R^+$ and R^+ is **transitive**:

$$(x, y) \in R^+ \wedge (y, z) \in R^+ \rightarrow (x, z) \in R^+.$$

Then

$$R^+ = \bigcup_{n \in \{1, 2, \dots\}} R^n.$$

- Happily, all of our relations are **finite**, so there must be some k with

$$R^+ = R \cup R^2 \cup \dots \cup R^k.$$

- Sadly, k will depend on the contents of R !
- Conclude: we **cannot** compute transitive closure in the Relational Algebra (or SQL without recursion).

ANOTHER CONTEST!!!

The challenge

HyperSQL implements SQL's **notoriously complicated** recursive query constructs. Write a recursive query that takes n as a parameter and computes all actors with Bacon number n . How large an n can your machine handle? On the small database? On the large database?

Send solutions to `tgg22@cam.ac.uk`.

PRIZES!!!



Observations

- We could continue, but the queries to compute higher and higher bacon numbers will grow in size and complexity.
- Performance will degrade rapidly.

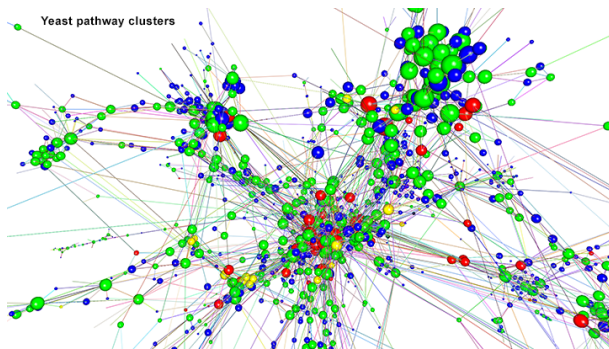
BIG IDEA

Is it possible to design and implement a database system that is **optimised for transitive closure and related path-oriented queries**?

Yes: This is one of the motivations behind graph-oriented databases.

We will be using Neo4j

- A Neo4j database contains **nodes** and binary **relationships** between nodes.
- Nodes and relationships can have attributes (called **properties**).
- Neo4j has a query language called **Cypher** that contains path-oriented constructs. It is designed to explore very large graphs.



Let's compute **all** Bacon numbers

```
MATCH (p:Person)
where p.name <> "Bacon, Kevin (I)"
with p
match paths=allshortestpaths(
    (m:Person {name : "Bacon, Kevin (I)"} )
    -[:ACTS_IN*]- (n:Person {name : p.name}))
return distinct p.name,
    length(paths)/2 as bacon_number
order by bacon_number desc;
```

On the small database this returns 5574 rows (out of 6458 actors).

p.name	bacon_number
"Alberson, Sarah"	5
"Allee, Virginia (I)"	5
"Allen, Brian Keith"	5
"Allen, Tess"	5
"Arbogast, Jeremiah"	5
....	...

This was computed in 10 seconds on my shabby laptop!

Let's compute count Bacon numbers

```
MATCH (p:Person)
where p.name <> "Bacon, Kevin (I)"
with p
match paths=allshortestpaths(
    (m:Person {name : "Bacon, Kevin (I)"} )
    -[:ACTS_IN*]- (n:Person {name : p.name}))
return count(distinct p.name) as total,
       length(paths)/2 as bacon_number
order by bacon_number desc;
```

+-----+		
total	bacon_number	
+-----+		
232	5	
668	4	
3186	3	
1387	2	
101	1	
+-----+		

Lecture 8 : Neo4j Practical Lecture

Before lecture, please attempt **Getting started with Neo4j** and **Neo4j exercises** found on the course web page.

Lecture 10

- Optimise for reading data?
- Document-oriented databases
- Semi-structured data
- Our bespoke database: DoctorWho
- Using Java as a query language

Optimise for reading

A fundamental tradeoff

Introducing data redundancy can speed up read-oriented transactions at the expense of slowing down write-oriented transactions.

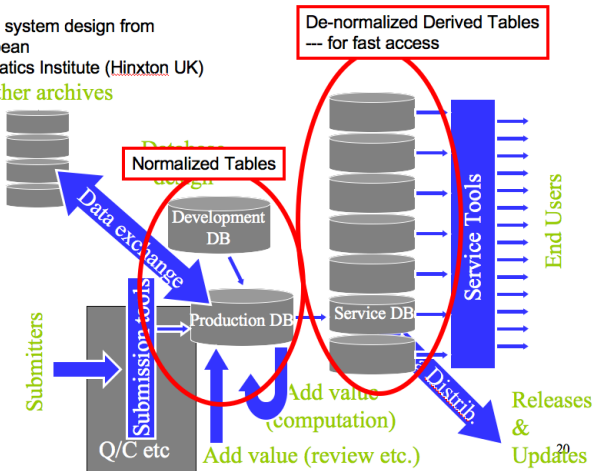
Situations where we might want a read-oriented database

- Your data is seldom updated, but very often read.
- Your reads can afford to be mildly out-of-synch with the write-oriented database. Then consider periodically extracting read-oriented snapshots and storing them in a database system optimised for reading. The following two slides illustrate examples of this situation.

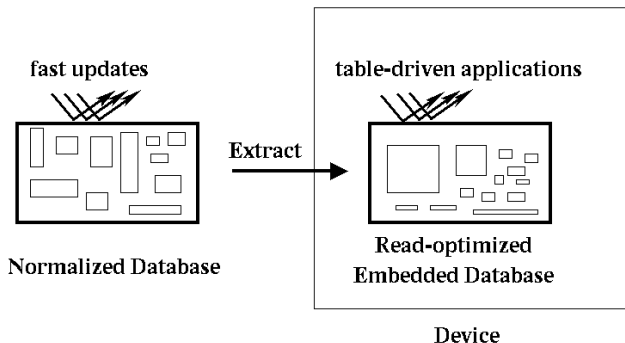
Example : Hinxton Bio-informatics

Database system design from
the European
Bioinformatics Institute (Hinxton UK)

Other archives



Example : Embedded databases



FIDO = Fetch Intensive Data Organization

Semi-structured data : JSON

```
{ "menu": {  
  "id": "file",  
  "value": "File",  
  "popup": {  
    "menuitem": [  
      { "value": "New", "onclick": "CreateNewDoc()" },  
      { "value": "Open", "onclick": "OpenDoc()" },  
      { "value": "Close", "onclick": "CloseDoc()" }  
    ]  
  }  
}}
```

From <http://json.org/example.html>.

Semi-structured data : XML

```
<menu id="file" value="File">
  <popup>
    <menuitem value="New" onclick="CreateNewDoc()" />
    <menuitem value="Open" onclick="OpenDoc()" />
    <menuitem value="Close" onclick="CloseDoc()" />
  </popup>
</menu>
```

From <http://json.org/example.html>.

Document-oriented database systems

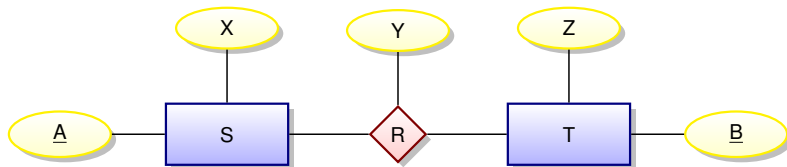
Our working definition

A **document-oriented databases** stores data in the form of **semi-structured objects**. Such database systems are also called **aggregate-oriented databases**.

Why Semi-structured data?

- Let's do **a thought experiment**.
- In the next few slides imagine that we intend to use a relational database to store read-optimised tables generated from a a set of write-optimised tables (that is, having little redundancy).
- We will encounter some problems that can be solved by representing our data as semi-structured objects.

Start with a simple relationship ...



A database instance

<i>S</i>	
<i>A</i>	<i>X</i>
a1	x1
a2	x2
a3	x3

<i>R</i>		
<i>A</i>	<i>B</i>	<i>Y</i>
a1	b1	y1
a1	b2	y2
a1	b3	y3
a2	b1	y4
a2	b3	y5

<i>T</i>	
<i>B</i>	<i>Z</i>
b1	z1
b2	z2
b3	z3
b4	z4

Imagine that our read-oriented applications can't afford to do joins!

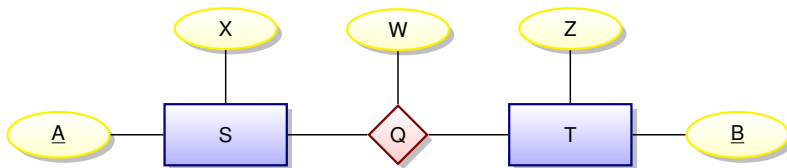
Implement the relationship as one big table?

BigTableOne: An outer join of S , R , and T

<u>A</u>	X	<u>B</u>	Z	Y
a1	x1	b1	z1	y1
a1	x1	b2	z2	y2
a1	x1	b3	z3	y3
a2	x2	b1	z1	y4
a2	x2	b3	z3	y5
a3	x3			
		b4	z4	

Since we don't update this data we will not encounter the problems associated with redundancy.

However, we might have many more relationships ...



A database instance

S	
A	X
a1	x1
a2	x2
a3	x3

Q		
A	B	W
a1	b4	w1
a3	b2	w2
a3	b3	w3

T	
B	Z
b1	z1
b2	z2
b3	z3
b4	z4

Implement with another big table?

BigTableTwo: An outer join of S , Q , and T

<u>A</u>	X	<u>B</u>	Z	W
a1	x1	b4	z4	w1
a3	x3	b2	z2	w2
a3	x3	b3	z3	w3
a2	x2			
		b1	z1	

Having two tables makes reading a bit more difficult!

Combine into one big table?

BigTable: Derived from S , R , Q , and T

<u>A</u>	X	<u>B</u>	Z	Y	W
a1	x1	b1	z1	y1	
a1	x1	b2	z2	y2	
a1	x1	b3	z3	y3	
a2	x2	b1	z1	y4	
a2	x2	b3	z3	y5	
a1	x1	b4	z4		w1
a3	x3	b2	z2		w2
a3	x3	b3	z3		w3

Problems with BigTable

- We could store BigTable and speed up some queries.
- But suppose that our applications typically access data using either S 's key or T 's key.
- Creating indices on the A and B columns could speed things up, but our applications may still be forced to gather information from many rows in order to collect all information related to a given key of S or a given key of T .
- It would be better to access all data associated with a given key of S or a given key of T using only **a single database lookup**.

Potential Solution

Represent the data using semi-structured objects.

Use (S-oriented) documents ("A" value is unique id)

```
{  "A": a1, "X": x1,  
   "R": [{ "B": b1, "Z": z1, "Y": y1 },  
         { "B": b2, "Z": z2, "Y": y2 },  
         { "B": b3, "Z": z3, "Y": y3 }],  
   "Q": [{ "B": b4, "Z": z4, "W": w1 }]  
}
```

```
{  "A": a2, "X": x2,  
   "R": [{ "B": b1, "Z": z1, "Y": y4 },  
         { "B": b3, "Z": z3, "Y": y5 }],  
   "Q": []  
}
```

```
{  "A": a3, "X": x3,  
   "R": [],  
   "Q": [{ "B": b2, "Z": z2, "W": w2 },  
         { "B": b3, "Z": z3, "W": w3 }]  
}
```


Use (*T*-oriented) documents ("B" value is id)

```
{ "B": b1, "Z": z1,  
  "R": [{ "A": a1, "X": x1, "Y": y2},  
         { "A": a2, "X": x2, "Y": y4}],  
  "Q": [] }  
  
{ "B": b2, "Z": z2,  
  "R": [{ "A": a1, "X": x1, "Y": y2}],  
  "Q": [{ "A": a3, "X": x3, "Y": w2}] }  
  
{ "B": b3, "Z": z3,  
  "R": [{ "A": a1, "X": x1, "Y": y3},  
         { "A": a2, "X": x2, "Y": y5}],  
  "Q": [{ "A": a3, "X": x3, "Y": w3}] }  
  
{ "B": b4, "Z": z4, "R": [],  
  "Q": [{ "A": a1, "X": x1, "Y": w1}] }
```

IMDb Person example

```
{ "id": 402542, "name": "Coen, Joel", "gender": "male",
  "director_in": [
    { "movie_id": 3256273, "title": "No Country for Old Men (2007)" },
    { "movie_id": 3667358, "title": "True Grit (2010)" },
    { "movie_id": 3026002, "title": "Inside Llewyn Davis (2013)"}],
  "editor_in": [
    { "movie_id": 3256273, "title": "No Country for Old Men (2007)" },
    { "movie_id": 3667358, "title": "True Grit (2010)" },
    { "movie_id": 3026002, "title": "Inside Llewyn Davis (2013)"}],
  "producer_in": [
    { "movie_id": 3256273, "title": "No Country for Old Men (2007)" },
    { "movie_id": 3667358, "title": "True Grit (2010)" },
    { "movie_id": 3026002, "title": "Inside Llewyn Davis (2013)"}],
  "writer_in": [
    { "line_order": 1, "group_order": 1, "subgroup_order": 1,
      "movie_id": 3256273, "title": "No Country for Old Men (2007)",
      "note": "(screenplay)" },
    { "line_order": 1, "group_order": 1, "subgroup_order": 1,
      "movie_id": 3667358, "title": "True Grit (2010)",
      "note": "(screenplay)" },
    { "line_order": 1, "group_order": 1, "subgroup_order": 1,
      "movie_id": 3026002, "title": "Inside Llewyn Davis (2013)",
      "note": "(written by)" }
  ]
}
```

IMDb Movie example (greatly simplified)

```
{
  "title": "Mad Max: Fury Road (2015)",
  "year": 2015,
  "actors": [
    {
      "character": "Max Rockatansky",
      "name": "Hardy, Tom (I)"
    },
    {
      "character": "Imperator Furiosa",
      "name": "Theron, Charlize"
    }
  ],
  "directors": [
    { "name": "Miller, George (II)" }
  ]
}
```

Key-value stores

- One of the simplest types of database systems is the **key-value store** that simply maps a key to a block of bytes.
- The retrieved block of bytes is typically opaque to the databases system.
- Interpretation of such data is left to applications.

This describes what might be called a **pure** key-value store. Some key-value stores extend this architecture with some limited capabilities to inspect blocks of data and extract meta-data such as indices. This is the case with **Berkeley DB** used to implement our bespoke data store **DoctorWho**.

How do we query DoctorWho?

We write code!

```
import uk.ac.cam.cl.databases.moviedb.MovieDB;
import uk.ac.cam.cl.databases.moviedb.model.*;

public class GetMovieById {
    public static void main(String[] args) {
        try (MovieDB database = MovieDB.open(args[0])) {
            int id = Integer.parseInt(args[1]);
            Movie movie = database.getMovieById(id);
            System.out.println(movie);
        }
    }
}
```

This code takes two command-line arguments: the directory containing the database and a movie id. It prints the associated JSON object.

Lecture 11 : DoctorWho Practical Lecture

Before lecture, please attempt **Getting started with DoctorWho** and **DoctorWho exercises** found on the course web page.

Lecture 12

- Sarah Mei's blog
- **OnLine Analytical Processing (OLAP)**
- **OnLine Transaction Processing (OLTP)**
- Cloud computing and distributed databases
- Column-oriented databases
- **Consistency, Availability, Partition tolerance (CAP)**

MongoDB is a popular document-oriented DBMS

Sarah Mei



« [What Your Conference Proposal Is Missing](#)

[How To Prevent Inappropriate Presentations](#)

Why You Should Never Use MongoDB

Disclaimer: I do not build database engines. I build web applications. I run 4-6 different projects every year, so I build **a lot** of web applications. I see apps with different requirements and different data storage needs. I've deployed most of the data stores you've heard about, and a few that you probably haven't.

Blog posted 11 November, 2013.

<http://www.sarahmei.com/blog/2013/11/11/why-you-should-never-use-mongodb>

Quotes from Sarah Mei's blog

- For quite a few years now, the received wisdom has been that **social data is not relational**, and that if you store it in a relational database, you're doing it wrong.
- Diaspora chose MongoDB for their social data in this zeitgeist. It was not an unreasonable choice at the time, given the information they had.
- You can see why this is attractive: all the data you need is already located where you need it.

Note

The blog describes implementation decisions made in 2010 for the development of a social networking platform called Diaspora. This was before mature graph-oriented databases were available.

Quotes from Sarah Mei's blog (2)

- You can also see why this is dangerous. Updating a user's data means walking through all the activity streams that they appear in to change the data in all those different places. This is very error-prone, and often leads to inconsistent data and mysterious errors, particularly when dealing with deletions.
- If your data looks like that, you've got documents. Congratulations! It's a good use case for Mongo. But if there's value in the links between documents, then you don't actually have documents. MongoDB is not the right solution for you. It's certainly not the right solution for social data, where links between documents are actually the most critical data in the system.

The project eventually migrated to a relational database (using about 50 tables). See <https://github.com/diaspora/diaspora>.

Yet another class of read-oriented databases

OLAP vs. OLTP

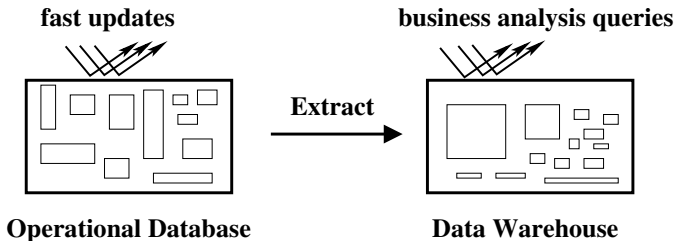
OLTP Online Transaction Processing

OLAP Online Analytical Processing

- Commonly associated with terms like Decision Support, Data Warehousing, etc.

	OLAP	OLTP
Supports	analysis	day-to-day operations
Data is	historical	current
Transactions mostly	reads	updates
optimized for	reads	updates
data redundancy	high	low
database size	humongous	large

Example : Data Warehouse (Decision support)



Limits of SQL aggregation

sale	prodlid	storeid	amt
	p1	c1	12
	p2	c1	11
	p1	c3	50
	p2	c2	8



	c1	c2	c3
p1	12		50
p2	11	8	

- Flat tables are great for processing, but hard for people to read and understand.
- Pivot tables and cross tabulations (spreadsheet terminology) are very useful for presenting data in ways that people can understand.
- Note that some table **values** become column or row **names**!
- Standard SQL does not handle pivot tables and cross tabulations well.

A very influential paper [G+1997]

Data Mining and Knowledge Discovery 1, 29–53 (1997)
© 1997 Kluwer Academic Publishers. Manufactured in The Netherlands.

Data Cube: A Relational Aggregation Operator Generalizing Group-By, Cross-Tab, and Sub-Totals*

JIM GRAY
SURAJIT CHAUDHURI
ADAM BOSWORTH
ANDREW LAYMAN
DON REICHART
MURALI VENKATRAO

Microsoft Research, Advanced Technology Division, Microsoft Corporation, One Microsoft Way, Redmond, WA 98052

Gray@Microsoft.com
SurajitC@Microsoft.com
AdamB@Microsoft.com
AndrewL@Microsoft.com
DonRei@Microsoft.com
MuraliV@Microsoft.com

FRANK PELLOW
HAMID PIRAHESH
IBM Research, 500 Harry Road, San Jose, CA 95120

Pellow@vnet.IBM.com
Pirahesh@Almaden.IBM.com

From aggregates to data cubes

Aggregate



Sum

Group By (with total)

By Color

RED
WHITE
BLUE

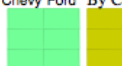


Sum

Cross Tab

Chevy Ford By Color

RED
WHITE
BLUE

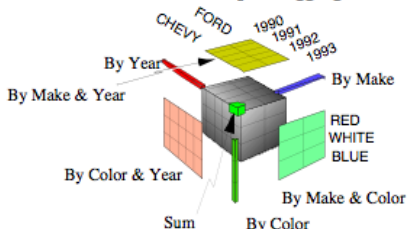


By Make

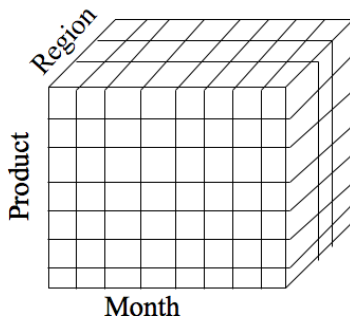


Sum

The Data Cube and The Sub-Space Aggregates



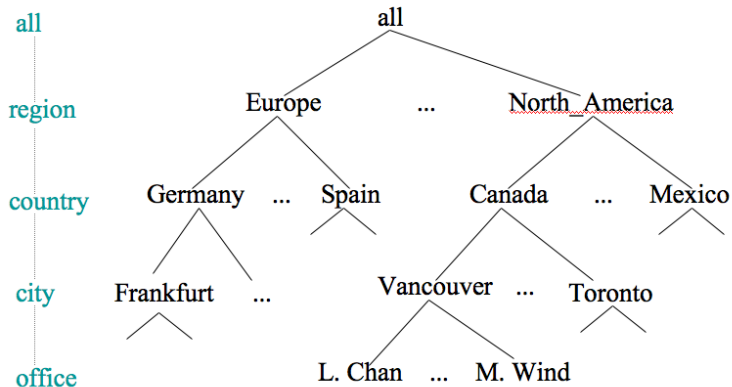
The Data Cube



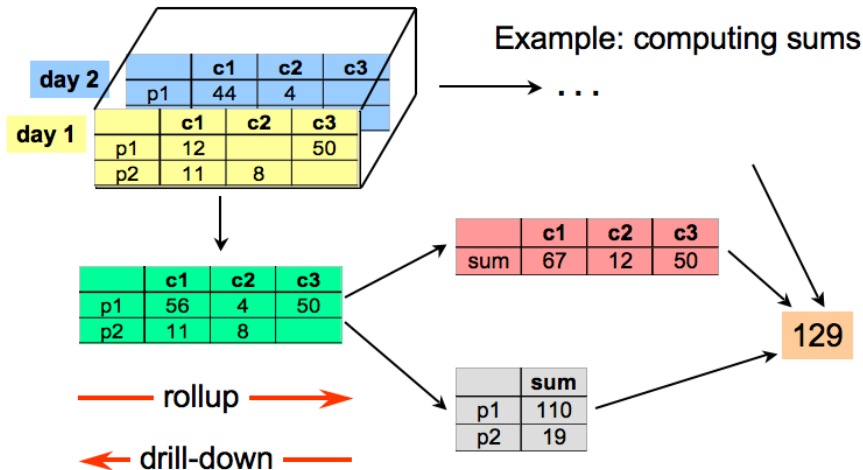
Dimensions:
Product,
Location,
Time

- Data modeled as an n -dimensional (hyper-) cube
- Each dimension is associated with a hierarchy
- Each “point” records facts
- Aggregation and cross-tabulation possible along all dimensions

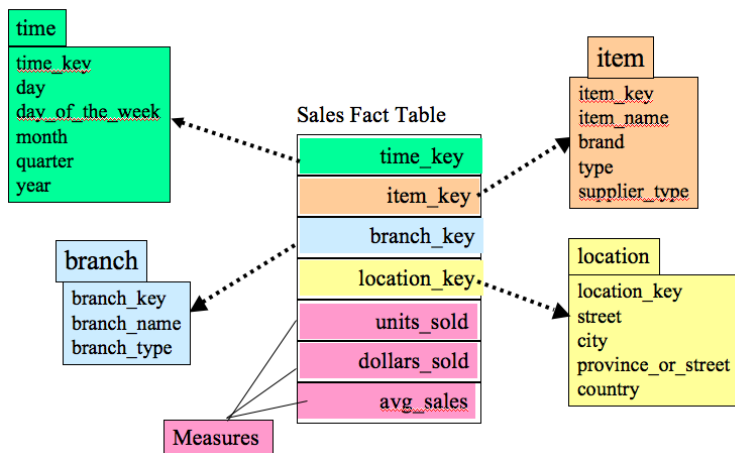
Hierarchy for **Location** Dimension



Cube Operations



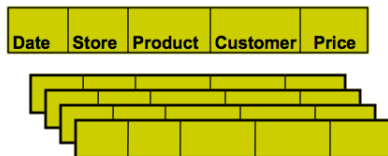
The Star Schema as a design tool



- In practice fact tables can be very large with hundreds of columns.
- Row-oriented table stores can be very inefficient since a typical query is concerned with only a few columns.

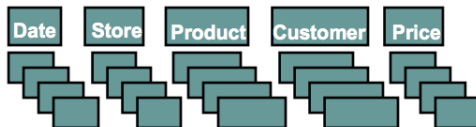
Column-oriented implementations

row-store



- + easy to add/modify a record
- might read in unnecessary data

column-store



- + only need to read in relevant data
- tuple writes require multiple accesses

=> suitable for read-mostly, read-intensive, large data repositories

From VLDB 2009 Tutorial: Column-Oriented Database Systems, by Stavros Harizopoulos, Daniel Abadi, Peter Boncz.

Distributed databases

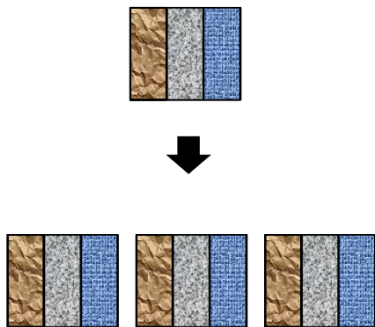
Why distribute data?

- **Scalability.** The data set or the workload can be too large for a single machine.
- **Fault tolerance.** The service can survive the failure of some machines.
- **Lower Latency.** Data can be located closer to widely distributed users.

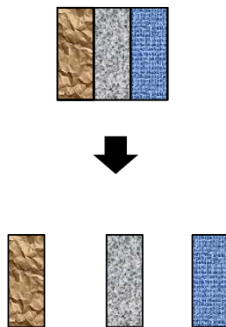
Distributed databases are an important technology supporting cloud computing.

How do we distribute the data?

Replication



Partitioning



Note: partitions themselves are often replicated.

Distributed databases pose difficult challenges

CAP concepts

- **Consistency.** All reads return data that is up-to-date.
- **Availability.** All clients can find some replica of the data.
- **Partition tolerance.** The system continues to operate despite arbitrary message loss or failure of part of the system.

It is very hard to achieve all three in a highly distributed database.

CAP principle

In a highly distributed system:

- Assume that network partitions and other connectivity problems will occur.
- Implementing transactional semantics is very difficult and slow.
- You are left engineering a **trade-off between availability and consistency**.

This gives rise to the notion of **eventual consistency**: if update activity ceases, then the system will eventually reach a consistent state.

What have we learned?

- Having a conceptual model of data is very useful, no matter which implementation technology is employed.
- There is a trade-off between fast reads and fast writes.
- There is no databases system that satisfies all possible requirements!
- It is best to understand pros and cons of each approach and develop integrated solutions where each component database is dedicated to doing what it does best.
- The future will see enormous churn and creative activity in the database field!

The End



(<http://xkcd.com/327>)