

Advanced Data Management Technologies

Unit 12 — Generalized Multidimensional Join

J. Gamper

Free University of Bozen-Bolzano
Faculty of Computer Science
IDSE

Acknowledgements: I am indebted to M. Böhlen for providing me the lecture notes.

Outline

- 1 Motivation
- 2 Definition of the GMD-Join
- 3 2D Cumulative Aggregates
- 4 Reduction to SQL
- 5 Computation of the GMD-Join
- 6 Experimental Evaluation

Outline

- 1 **Motivation**
- 2 Definition of the GMD-Join
- 3 2D Cumulative Aggregates
- 4 Reduction to SQL
- 5 Computation of the GMD-Join
- 6 Experimental Evaluation

Goals

- We need an algebraic operator for **Ad Hoc OLAP** queries (i.e., no data cube is pre-computed and aggregates are computed on the detail data):
 - The operator must allow us to **easily express** OLAP queries.
 - Grouping should be independent from detail data.
 - More general specification of groups over which aggregates are computed.
 - The operator must admit **efficient evaluation** algorithms.
 - The operator must **interact seamlessly** with the other operators of the relational algebra (query optimization).

Motivating Application

- Analysis of network data
 - Collect, correlate, and analyze huge amounts of data across the network.
- Example queries:
 - **Network usage:** For each IP address, what fraction of the total traffic is due to web flows?
 - **Principal components:** On an hourly basis, what fraction of the total traffic is from IP subnets whose hourly traffic is within 10% of the max?
 - **Pattern identification:** Break down all flows recorded on US election day by all combinations of source router, destination router, and protocol.

Data Warehouse Schema

- IP Flows star schema

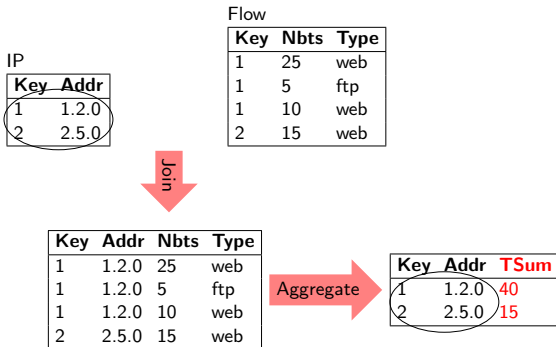
- UserDim(UserId, Name, Addr, Balance, ...)
- IPRegDim(IPKey, UserId, IP, Port, Mask, ASsystem, ...)
- RouterDim(RouterId, Addr, Type, ...)
- HourDim(HourKey, Start, End)
- FlowFact(RouterId, SourceKey, DestKey, Protocol, StartTime, EndTime, NumPackets, NumBytes, ...)

- Denormalized IP Flow schema

- Flow(RouterId, SourceIP, SourcePort, SourceAS, DestIP, DestPort, DestAS, Protocol, StartTime, EndTime, NumPackets, NumBytes, ...)

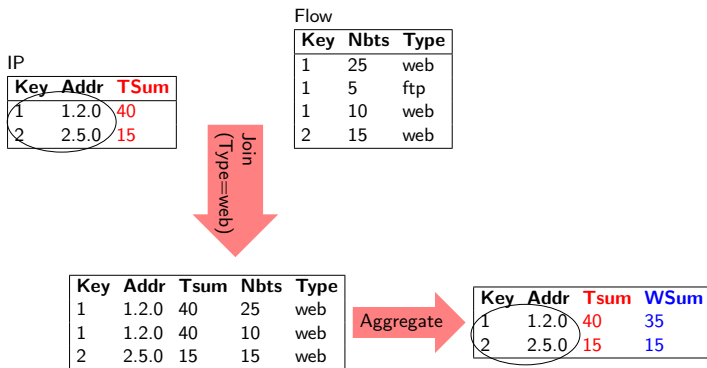
Network Usage (Relational)/1

- **Network usage:** For each IP address, what fraction of the total traffic is due to web flows?
- Step 1: Determine **total traffic**.



Network Usage (Relational)/2

- Step 2: Determine **web** traffic.



- Combination of joins and aggregations with different groupings are required.
- OLAP extensions provide some support, but some queries are still difficult to express and might be expensive.

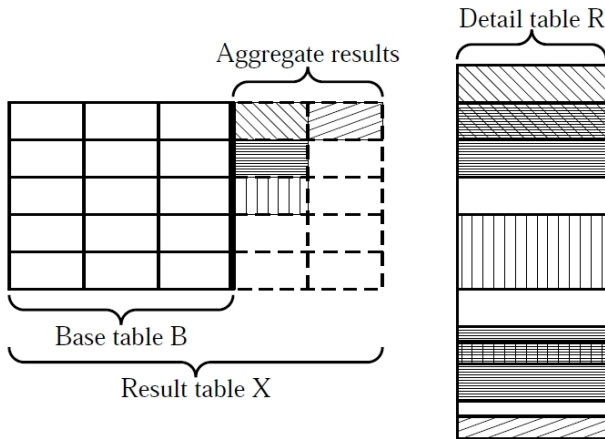
Outline

- 1 Motivation
- 2 Definition of the GMD-Join**
- 3 2D Cumulative Aggregates
- 4 Reduction to SQL
- 5 Computation of the GMD-Join
- 6 Experimental Evaluation

The Generalized MD-join/1

- $\text{MD}(\mathbf{b}, \mathbf{r}, (l_1, \dots, l_m), (\theta_1, \dots, \theta_m))$
 - $\mathbf{b}$ is the base table (the “groups”);
 - $\mathbf{r}$ is the detail table (fact data);
 - l_1 to l_m are lists of aggregate functions;
 - θ_1 to θ_m are possibly complex conditions over $\mathbf{b}$ and $\mathbf{r}$ describing what data in $\mathbf{r}$ is to be aggregated for each l_i , respectively.
- Result: base table $\mathbf{b}$ **extended** with the aggregates in l_1 to l_m .
- Salient feature: decouples grouping from aggregation
 - SQL does not do this!

The Generalized MD-join/2



Network Usage (GMD-join)/1

- **Network usage:** For each IP address, what fraction of the total number of flows is due to web traffic?
- Step 1: total traffic
 - $\mathbf{x}_1 = \text{MD}(\text{IP}, \text{Flow}, (\text{SUM}(\text{NBts}) / \text{TSum}), (\text{IP}.\text{Key} = \text{Flow}.\text{Key}))$
 - Schema of $\mathbf{x}_1$: (Key, Addr, TSum)
- Step 2: web traffic
 - $\mathbf{x}_2 = \text{MD}(\mathbf{x}_1, \text{Flow}, (\text{SUM}(\text{NBts}) / \text{WSum}), (\mathbf{x}_1.\text{Key} = \text{Flow}.\text{Key} \wedge \text{Flow}.\text{Source} = \text{'web'}))$
 - Schema of $\mathbf{x}_2$: (Key, Addr, TSum, WSum)
- Sequences of GMDJs instead of multiple aggregate-join expressions.

Network Usage (GMD-join)/2

- GMD-join can do it in one step
 - MD(IP, Flow, ((SUM(NBts)/ TSum), (SUM(NBts)/ WSum)), (θ_1, θ_2))
 - θ_1 : (IP.Key = Flow.Key)
 - θ_2 : (IP.Key = Flow.Key $\wedge$ Flow.Type = ' web')
- Equivalent RA expression (left outer join, LOJ):

$$\pi[\text{Key}, \text{IP}, \text{S1}, \text{SUM}(\text{NB})]($$

$$(\pi[\text{Key}, \text{IP}, \text{SUM}(\text{NB})/\text{S1}](\text{IP LOJ}[\theta_1] \text{ Flow}))$$

$$\text{LOJ}[\theta_2] \text{ Flow})$$

- θ_1 : (IP.Key = Flow.Key)
- θ_2 : (IP.Key = Flow.Key $\wedge$ Flow.Type = ' web')

Execution of GMD-join/1

- MD(IP, Flow, ($\text{SUM}(\text{NBts}) / \text{TSum}$, $\text{SUM}(\text{NBts}) / \text{WSum}$), (θ_1, θ_2))
 - θ_1 : (IP.Key = Flow.Key)
 - θ_2 : (IP.Key = Flow.Key $\wedge$ Flow.Type = web)
- Step 1: Init base-result structure

IP		Result table x			
Key	Addr				
1	1.2.0				
2	2.5.0				

 $\xrightarrow{\text{Init}}$

Key	Addr	TSum	WSum
1	1.2.0	0	0
2	2.5.0	0	0

- Step 2: Scan detail tuples and incrementally update result structure.

Result table x				Flow		
Key	Addr	TSum	WSum			
1	1.2.0	25	25			
2	2.5.0	0	0			

 $\xleftarrow{\text{Update}}$

Key	Nbts	Type
1	25	web
1	5	ftp
1	10	web
2	15	web

Execution of GMD-join/2

Result table x

Key	Addr	TSum	WSum
1	1.2.0	30	25
2	2.5.0	0	0

Flow

Key	Nbts	Type
1	25	web
1	5	ftp
1	10	web
2	15	web

Update
←

Result table x

Key	Addr	TSum	WSum
1	1.2.0	40	35
2	2.5.0	0	0

Flow

Key	Nbts	Type
1	25	web
1	5	ftp
1	10	web
2	15	web

Update
←

Result table x

Key	Addr	TSum	WSum
1	1.2.0	40	35
2	2.5.0	15	15

Flow

Key	Nbts	Type
1	25	web
1	5	ftp
1	10	web
2	15	web

Update
←

Advantages of the GMD-join

- **Flexible** Operator that allows a “simple” formulation of very **complex** queries
 - Evaluation of multiple complex aggregations using just a single operator and a **single pass** through the detail data.
- Very **efficient** to evaluate
 - Indexing on the base-result structure constructed at query execution time performs very effectively.
- Can be **optimized** effectively
 - GMD-joins allow efficient optimization schemes to be developed for complex aggregate queries.

Complex Aggregation Example/1

- Assume the following denormalized IP Flow schema:
Flow(RouterId, SourceIP, SourcePort, SourceAS, DestIP, DestPort, DestAS, Protocol, StartTime, EndTime, NumPackets, NumBytes, ...)
- **Query:** Determine the **cumulative total**, **two hour moving average**, and **hourly total** of network flow broken down by hour of the day.

Complex Aggregation Example/2

- GMD-Join:

MD(Hour/**b**, Flow/**r**, (l_1, l_2, l_3), ($\theta_1, \theta_2, \theta_3$))

- $l_1 : (SUM(NB)/CSum)$
- $\theta_1 : (r.T \leq b.HEnd)$
- $l_2 : (SUM(NB)/MSum, COUNT(*)/Mcnt)$
- $\theta_2 : (r.T \geq b.HStart - 60 \wedge r.T \leq b.HEnd)$
- $l_3 : (SUM(NB)/Hsum)$
- $\theta_3 : (r.T \geq b.HStart \wedge r.T \leq b.HEnd)$

- Result table

HId	HStart	HEnd	CSum	MAvg	HSum
1	0	59	11	11/2	11
2	60	119	17	17/3	6
3	120	179	33	22/3	16

- Base and detail table

Hour

HId	HStart	HEnd
1	0	59
2	60	119
3	120	179

Flow

SIP	DIP	NB	T
5	29	3	30
5	7	8	45
5	29	6	110
7	29	6	161
6	29	10	170

- Can be expressed in SQL with window functions!

Outline

- 1 Motivation
- 2 Definition of the GMD-Join
- 3 2D Cumulative Aggregates**
- 4 Reduction to SQL
- 5 Computation of the GMD-Join
- 6 Experimental Evaluation

2D Cumulative Aggregates Example/1

- Consider the following *Lineitem* relation

Lineitem

	<i>Ordkey</i>	<i>Partkey</i>	<i>Suppkey</i>	<i>Quantity</i>	<i>Price</i>	<i>Discount</i>	<i>Shipdate</i>
r_1	O1	P1	S1	2	220	0.00	2008.01.23
r_2	O2	P1	S1	4	440	0.05	2008.01.23
r_3	O3	P2	S1	6	300	0.10	2008.01.23
r_4	O4	P2	S2	7	420	0.10	2008.01.23
r_5	O5	P2	S1	2	100	0.00	2008.01.24
r_6	O6	P1	S2	3	240	0.05	2008.01.24
r_7	O7	P2	S1	9	450	0.05	2008.01.24
r_8	O8	P1	S2	8	640	0.10	2008.01.24

- Query Q1:** Compute the number of orders per day and disc rate (*CntDD*), the cumulative number of orders per day (*CumCntD*), and the cumulative number of orders per day and disc rate (*CumCntDD*).

2D Cumulative Aggregates Example/2

<i>Lineitem</i>							
	<i>Ordkey</i>	<i>Partkey</i>	<i>Suppkey</i>	<i>Quantity</i>	<i>Price</i>	<i>Discount</i>	<i>Shipdate</i>
r_1	O1	P1	S1	2	220	0.00	2008.01.23
r_2	O2	P1	S1	4	440	0.05	2008.01.23
r_3	O3	P2	S1	6	300	0.10	2008.01.23
r_4	O4	P2	S2	7	420	0.10	2008.01.23
r_5	O5	P2	S1	2	100	0.00	2008.01.24
r_6	O6	P1	S2	3	240	0.05	2008.01.24
r_7	O7	P2	S1	9	450	0.05	2008.01.24
r_8	O8	P1	S2	8	640	0.10	2008.01.24

- Result x of Q1

x					
	<i>Shipdate</i>	<i>Discount</i>	<i>CntDD</i>	<i>CumCntD</i>	<i>CumCntDD</i>
x_1	2008.01.23	0.00	1	4	1
x_2	2008.01.23	0.05	1	4	2
x_3	2008.01.23	0.10	2	4	4
x_4	2008.01.24	0.00	1	8	2
x_5	2008.01.24	0.05	2	8	5
x_6	2008.01.24	0.10	1	8	8

2D Cumulative Aggregates Example/3

- GMD-join formulation

$$\text{MD}(\pi[\text{Shipdate}, \text{Discount}]\text{Lineitem}/\mathbf{b}, \text{Lineitem}/\mathbf{r}, (l_1, l_2, l_3), (\theta_1, \theta_2, \theta_3))$$

$$l_1 \equiv (\text{COUNT}(\text{Quantity})/\text{CntDD})$$

$$\theta_1 \equiv (\mathbf{r}.\text{Shipdate} = \mathbf{b}.\text{Shipdate} \wedge \mathbf{r}.\text{Discount} = \mathbf{b}.\text{Discount})$$

$$l_2 \equiv (\text{COUNT}(\text{Quantity})/\text{CumCntD})$$

$$\theta_2 \equiv (\mathbf{r}.\text{Shipdate} \leq \mathbf{b}.\text{Shipdate})$$

$$l_3 \equiv (\text{COUNT}(\text{Quantity})/\text{CumCntDD})$$

$$\theta_3 \equiv (\mathbf{r}.\text{Shipdate} \leq \mathbf{b}.\text{Shipdate} \wedge \mathbf{r}.\text{Discount} \leq \mathbf{b}.\text{Discount})$$

2D Aggregates (*CntDD*)

- Groups are defined by **identical values** on the grouping attributes.
- Can be expressed in SQL
 - GROUP BY *Shipdate*, *Discount*

Lineitem/r

	<i>Ordkey</i>	<i>Partkey</i>	<i>Suppkey</i>	<i>Quantity</i>	<i>Price</i>	<i>Discount</i>	<i>Shipdate</i>
<i>r</i> ₁	O1	P1	S1	2	220	0.00	2008.01.23
<i>r</i> ₂	O2	P1	S1	4	440	0.05	2008.01.23
<i>r</i> ₃	O3	P2	S1	6	300	0.10	2008.01.23
<i>r</i> ₄	O4	P2	S2	7	420	0.10	2008.01.23
<i>r</i> ₅	O5	P2	S1	2	100	0.00	2008.01.24
<i>r</i> ₆	O6	P1	S2	3	240	0.05	2008.01.24
<i>r</i> ₇	O7	P2	S1	9	450	0.05	2008.01.24
<i>r</i> ₈	O8	P1	S2	8	640	0.10	2008.01.24

$r.Shipdate = x.Shipdate \wedge$
 $r.Discount = x.Discount$

x

	<i>Shipdate</i>	<i>Discount</i>	<i>CntDD</i>	<i>CumCntD</i>	<i>CumCntDD</i>
<i>x</i> ₁	2008.01.23	0.00	1	4	1
<i>x</i> ₂	2008.01.23	0.05	1	4	2
<i>x</i> ₃	2008.01.23	0.10	2	4	4
<i>x</i> ₄	2008.01.24	0.00	1	8	2
<i>x</i> ₅	2008.01.24	0.05	2	8	5
<i>x</i> ₆	2008.01.24	0.10	1	8	8

1D Cumulative Aggregates (*CumCntD*)

- Groups are **contiguous rows** after sorting the data on the grouping attribute.
- Can be expressed by SQL window functions
 - `SUM(COUNT(*)) OVER (ORDER BY Shipdate ROWS UNBOUNDED PRECEDING)`

Lineitem/r

	Ordkey	Partkey	Suppkey	Quantity	Price	Discount	Shipdate
<i>r</i> ₁	O1	P1	S1	2	220	0.00	2008.01.23
<i>r</i> ₂	O2	P1	S1	4	440	0.05	2008.01.23
<i>r</i> ₃	O3	P2	S1	6	300	0.10	2008.01.23
<i>r</i> ₄	O4	P2	S2	7	420	0.10	2008.01.23
<i>r</i> ₅	O5	P2	S1	2	100	0.00	2008.01.24
<i>r</i> ₆	O6	P1	S2	3	240	0.05	2008.01.24
<i>r</i> ₇	O7	P2	S1	9	450	0.05	2008.01.24
<i>r</i> ₈	O8	P1	S2	8	640	0.10	2008.01.24

$r.Shipdate \leq x.Shipdate$

	Shipdate	Discount	CntDD	CumCntD	CumCntDD
<i>x</i> ₁	2008.01.23	0.00	1	4	1
<i>x</i> ₂	2008.01.23	0.05	1	4	2
<i>x</i> ₃	2008.01.23	0.10	2	4	4
<i>x</i> ₄	2008.01.24	0.00	1	8	2
<i>x</i> ₅	2008.01.24	0.05	2	8	5
<i>x</i> ₆	2008.01.24	0.10	1	8	8

2D Cumulative Aggregates (*CumCntDD*)

- Groups are specified along 2 dimensions
- Ordering with **contiguous rows** might **not** exist.

Lineitem/r

	Ordkey	Partkey	Suppkey	Quantity	Price	Discount	Shipdate
r ₁	O1	P1	S1	2	220	0.00	2008.01.23
r ₂	O2	P1	S1	4	440	0.05	2008.01.23
r ₃	O3	P2	S1	6	300	0.10	2008.01.23
r ₄	O4	P2	S2	7	420	0.10	2008.01.23
r ₅	O5	P2	S1	2	100	0.00	2008.01.24
r ₆	O6	P1	S2	3	240	0.05	2008.01.24
r ₇	O7	P2	S1	9	450	0.05	2008.01.24
r ₈	O8	P1	S2	8	640	0.10	2008.01.24

$r.\text{Shipdate} \leq x.\text{Shipdate} \wedge$
 $r.\text{Discount} \leq x.\text{Discount}$

x

	Shipdate	Discount	CntDD	CumCntD	CumCntDD
x ₁	2008.01.23	0.00	1	4	1
x ₂	2008.01.23	0.05	1	4	2
x ₃	2008.01.23	0.10	2	4	4
x ₄	2008.01.24	0.00	1	8	2
x ₅	2008.01.24	0.05	2	8	5
x ₆	2008.01.24	0.10	1	8	8

- 2D cumulative (and higher dimensional) aggregates might be **difficult and very expensive** in SQL!

Outline

- 1 Motivation
- 2 Definition of the GMD-Join
- 3 2D Cumulative Aggregates
- 4 Reduction to SQL**
- 5 Computation of the GMD-Join
- 6 Experimental Evaluation

Reduction to SQL

- The GMD-join $MD(\mathbf{b}, \mathbf{r}, (l_1, \dots, l_m), (\theta_1, \dots, \theta_m))$ can be systematically transformed into SQL as follows

```

SELECT  B1, ..., Bh
        f11 (CASE WHEN  $\theta_1$  THEN  $\mathbf{r}.A_{11}$  ELSE  $N_{11}$  END) AS C11
        ...
        fmk (CASE WHEN  $\theta_m$  THEN  $\mathbf{r}.A_{mk}$  ELSE  $N_{mk}$  END) AS Cmk
FROM     $\mathbf{b}$  LEFT OUTER JOIN  $\mathbf{r}$  ON  $\theta'$ 
GROUP BY B1, ..., Bh

```

- $B_1, \dots, B_h$ are the attributes of $\mathbf{b}$.
- f_{ij} are the aggregate functions in l_i .
- N_{ij} are the neutral values of aggregate functions.
- θ' is the common subpart of all θ_i
 - Cartesian Product if θ' is empty

Reduction to SQL Example/1

```
SELECT  B.Shipdate, B.Disc,  
        COUNT(CASE WHEN L.Shipdate = B.Shipdate AND L.Disc = B.Disc)  
            THEN Quantity  
            ELSE 0 END) AS CntDD,  
        COUNT(CASE WHEN L.Shipdate <= B.Shipdate)  
            THEN Quantity  
            ELSE 0 END) AS CumCntD,  
        COUNT(CASE WHEN L.Shipdate <= B.Shipdate AND L.Disc <= B.Disc)  
            THEN Quantity  
            ELSE 0 END) AS CumCntDD  
FROM    (SELECT DISTINCT Shipdate, Disc FROM Lineitem) B, Lineitem L  
GROUP BY B.Shipdate, B.Disc;
```

- Common condition $\theta' = \text{true}$ (empty), hence we have a Cartesian product.
- Efficient join techniques (e.g., hash or sort-merge joins) are not applicable.

Reduction to SQL Example/2

- Manual ad hoc optimizations are possible, e.g., compute all aggregates separately.

```
WITH
q1 AS ( SELECT   Shipdate, Disc, COUNT(Quantity) AS CntDD
        FROM     Lineitem
        GROUP BY Shipdate, Disc ),
q2 AS ( SELECT   Shipdate,
                  SUM(COUNT(Quantity)) OVER
                    (ORDER BY Shipdate ROWS UNBOUNDED PRECEDING) AS CumCntD
        FROM     Lineitem
        GROUP BY Shipdate, Disc ),
q3 AS ( SELECT   B.Shipdate, B.Disc, COUNT(Quantity) AS CumCntDD
        FROM     ( SELECT DISTINCT Shipdate, Disc FROM Lineitem ) B
        LEFT OUTER JOIN
        ( SELECT DISTINCT Shipdate, Disc FROM Lineitem ) L
        ON L.Shipdate <= B.Shipdate AND L.Disc <= B.Disc
        GROUP BY Shipdate, Disc ),
SELECT * FROM q1 NATURAL JOIN q2 NATURAL JOIN q3;
```

- 3 scans of base and detail table are required.

Outline

- 1 Motivation
- 2 Definition of the GMD-Join
- 3 2D Cumulative Aggregates
- 4 Reduction to SQL
- 5 Computation of the GMD-Join**
- 6 Experimental Evaluation

Basic Algorithm for the GMD-join

Algorithm: GMDJBasic($\mathbf{b}, \mathbf{r}, (l_1, \dots, l_m), (\theta_1, \dots, \theta_m)$)

// Initialize result table

$\mathbf{x} \leftarrow \mathbf{b}$ extended with a column for each aggregate;

Initialize $\mathbf{x}$ to NULL for MAX and MIN, to 0 for SUM and COUNT;

// Scan detail table and update aggregates

foreach $r \in \mathbf{r}$ **do**

foreach $x \in \mathbf{x}$ **do**

foreach $\theta_i \in \{\theta_1, \dots, \theta_m\}$ **do**

if $\theta_i(x, r)$ **then**

 Update aggregates l_i in x ;

return $\mathbf{x}$;

Hash Index for the GMD-join

Algorithm: GMDJHash($\mathbf{b}, \mathbf{r}, (l_1, \dots, l_m), (\theta_1, \dots, \theta_m)$)

// Initialize result table

$\mathbf{x} \leftarrow \mathbf{b}$ extended with a column for each aggregate;

Initialize $\mathbf{x}$ to NULL for MAX and MIN, to 0 for SUM and COUNT;

// Create an index on the result table

Construct hash index for $\mathbf{x}$;

// Scan detail table and update aggregates

foreach $r \in \mathbf{r}$ **do**

foreach $\theta_i \in \{\theta_1, \dots, \theta_m\}$ **do**

 Use index to find $\mathbf{x}_i \leftarrow \{x \mid x \in \mathbf{x} \wedge \theta_i(x, r)\}$;

foreach $x \in \mathbf{x}_i$ **do**

 Update aggregates l_i in x ;

return $\mathbf{x}$;

Incremental Aggregation/1

- The GMDJ **incrementally** computes the aggregate values in the result table while scanning the detail table.
- Different types of aggregate functions (AF) need to be distinguished.
 - Distributive AF
 - Algebraic AF
 - Holistic AF

Incremental Aggregation/2

• Distributive AF

- Result over a set A is the aggregation of partial results over subsets A_i
 - $A_1 \cup \dots \cup A_k = A, A_i \cap A_j = \emptyset$
- $F(A) = G(F(A_1), \dots, F(A_k))$
- G is also called **super-aggregate**
- Distributive aggregates: MIN, MAX, SUM, COUNT
 - e.g., $COUNT(A) = SUM(COUNT(A_1), \dots, COUNT(A_k))$

• Algebraic AF

- Computation can be reduced to a number of **distributive sub-aggregates**
- $F(A) = G(F_1(A), \dots, F_m(A))$
- Algebraic aggregates: AVG
 - $AVG(A) = DIV(SUM(A), COUNT(A))$

• Holistic AF

- Cannot be reduced to a fixed number of sub-aggregates.
- Incremental computation in GMD-join is **not possible**.
- $F(A) = F(A_1 \cup \dots \cup A_k)$
- Holistic aggregates: MEDIAN, MODE, RANK

Evaluation of Algebraic AF

Algorithm: GMDJAlgebraic($\mathbf{b}, \mathbf{r}, (l_1, \dots, l_m), (\theta_1, \dots, \theta_m)$)

// Replace algebraic aggregates by distributive sub-aggregates

Let $l'_i \leftarrow l_i$ for $1 \leq i \leq m$;

foreach algebraic aggregate $f_{ij} \in \{l'_1, \dots, l'_m\}$ **do**

 Replace f_{ij} with its distributive sub-aggregates $f_{ij}^1, \dots, f_{ij}^{p_{ij}}$;

// Compute aggregates

$\mathbf{x} \leftarrow \text{GMDJHash}(\mathbf{b}, \mathbf{r}, (l'_1, \dots, l'_m), (\theta_1, \dots, \theta_m));$

// Apply the super-aggregates

foreach $x \in \mathbf{x}$ **do**

foreach algebraic function $f_{ij} \in \{l_1, \dots, l_m\}$ **do**

 Let g_{ij} be the super-aggregate of f_{ij} ;

 In $\mathbf{x}$, replace columns $f_{ij}^1, \dots, f_{ij}^{p_{ij}}$ by a $g_{ij}(x.f_{ij}^1, \dots, f_{ij}^{p_{ij}})$;

return $\mathbf{x}$;

Algebraic AF Example/1

- For each combination of source (S) and destination (D), compute the number of flows whose NumBytes (NB) value exceeds the average value of NumBytes.
- Step 1: $\mathbf{x}_1 = \text{MD}(\pi[S, D]\mathbf{r}/\mathbf{b}, \mathbf{r}, l_1, \theta_1)$

$$l_1 = (\text{AVG}(\text{NB})/\text{Avg})$$

$$\theta_1 = (\mathbf{r}.S = \mathbf{b}.S \wedge \mathbf{r}.D = \mathbf{b}.D)$$
- l_1 contains algebraic aggregate AVG, hence it is translated into
 - $l_1 = (\text{COUNT}(*)/\text{Cnt1}, \text{SUM}(\text{NB})/\text{Sum1})$

Detail table **r**

S	D	NB
5	29	3
5	7	8
5	29	6
7	29	4
7	29	6
6	29	10

$\Rightarrow$

$\mathbf{x}_1$			
S	D	Cnt1	Sum1
5	29	2	9
5	7	1	8
7	29	2	10
6	29	1	10

$\Rightarrow$

$\mathbf{x}_1$		
S	D	Avg
5	29	4.5
5	7	8
7	29	5
6	29	10

Algebraic AF Example/2

- Step 2: $\mathbf{x} = \text{MD}(\mathbf{x}_1/\mathbf{b}, \mathbf{r}, l_2, \theta_2)$

$$l_2 = (\text{COUNT}(\text{NB})/\text{cnt2})$$

$$\theta_2 = (\mathbf{r}.S = \mathbf{b}.S \wedge \mathbf{r}.D = \mathbf{b}.D \wedge \mathbf{r}.NB > \mathbf{b}.Avg)$$

Detail table r

S	D	NB
5	29	3
5	7	8
5	29	6
7	29	4
7	29	6
6	29	10



x

S	D	Avg	Cnt2
5	29	4.5	1
5	7	8	0
7	29	5	1
6	29	10	0

Transformation Rules

- The GMD-Join interacts with the other operators of the relational algebra.

$$\pi[\mathbf{A}]\mathcal{G}^\theta(B, R, \vec{I}, \vec{\theta}) = \mathcal{G}^\theta(\pi[\mathbf{A}']B, R, \vec{I}, \vec{\theta})$$

$$\pi[\mathbf{A}]\mathcal{G}^\theta(B, R, \vec{I}, \vec{\theta}) = \pi[\mathbf{A}]\mathcal{G}^\theta(\pi[\mathbf{A}']B, R, \vec{I}, \vec{\theta})$$

$$\sigma[\theta_S]\mathcal{G}^\theta(B, R, \vec{I}, \vec{\theta}) = \mathcal{G}^\theta(\sigma[\theta_S]B, R, \vec{I}, \vec{\theta})$$

$$\mathcal{G}^\theta(B, R, \vec{I}, \vec{\theta}) = \mathcal{G}^\theta(B, \sigma[\theta_1^R \vee \dots \vee \theta_m^R]R, \vec{I}, \vec{\theta})$$

$$\sigma[>0]\mathcal{G}^\theta(B, R, \vec{I}, \vec{\theta}) = \sigma[>0]\mathcal{G}^\theta(\sigma[\theta_1^B \vee \dots \vee \theta_m^B]B, R, \vec{I}, \vec{\theta})$$

$$\mathcal{G}^\theta(\mathcal{G}^\theta(B, R_1, \vec{I}_1, \vec{\theta}_1), R_2, \vec{I}_2, \vec{\theta}_2) = \mathcal{G}^\theta(\mathcal{G}^\theta(B, R_2, \vec{I}_2, \vec{\theta}_2), R_1, \vec{I}_1, \vec{\theta}_1)$$

$$\mathcal{G}^\theta(\mathcal{G}^\theta(B, R, \vec{I}_1, \vec{\theta}_1), R, \vec{I}_2, \vec{\theta}_2) = \mathcal{G}^\theta(B, R, (\vec{I}_1, \vec{I}_2), (\vec{\theta}_1, \vec{\theta}_2))$$

$$\mathcal{G}^\theta(\mathcal{G}^\theta(B, R_1, \vec{I}_1, \vec{\theta}_1), R_2, \vec{I}_2, \vec{\theta}_2) = \mathcal{G}^\theta(B, R_1, \vec{I}_1, \vec{\theta}_1) \rightarrow_{\theta_B} \mathcal{G}^\theta(B, R_2, \vec{I}_2, \vec{\theta}_2) \rightarrow_V$$

$$\mathcal{G}^\theta(B, R, \vec{I}, \vec{\theta}) = \mathcal{G}^\theta(B_1, R, \vec{I}, \vec{\theta}) \cup \dots \cup \mathcal{G}^\theta(B_n, R, \vec{I}, \vec{\theta})$$

$$(attr(\vec{\theta}) \cap \mathbf{B}) \subseteq \mathbf{A}, \mathbf{A}' = \mathbf{A} \setminus \mathbf{C} \quad (\text{E1})$$

$$\mathbf{A}' = (\mathbf{A} \cup (attr(\vec{\theta}) \cap \mathbf{B})) \setminus \mathbf{C} \quad (\text{E2})$$

$$attr(\theta_S) \subseteq \mathbf{B} \quad (\text{E3})$$

$$\theta_i = \theta'_i \wedge \theta_i^R, attr(\theta_i^R) \subseteq \mathbf{R} \quad (\text{E4})$$

$$\theta_i = \theta'_i \wedge \theta_i^B, attr(\theta_i^B) \subseteq \mathbf{B} \quad (\text{E5})$$

$$attr(\vec{\theta}_2) \subseteq (\mathbf{B} \cup \mathbf{R}_2) \quad (\text{E6})$$

$$attr(\vec{\theta}_2) \subseteq (\mathbf{B} \cup \mathbf{R}) \quad (\text{E7})$$

$$attr(\vec{\theta}_2) \subseteq (\mathbf{B} \cup \mathbf{R}_2), \quad (\text{E8})$$

$$\theta_B = (U.B = V.B)$$

$$B = B_1 \cup \dots \cup B_n, \quad (\text{E9})$$

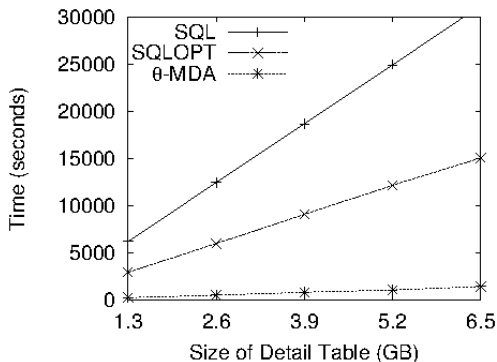
$$B_i \cap B_j = \emptyset$$

Outline

- 1 Motivation
- 2 Definition of the GMD-Join
- 3 2D Cumulative Aggregates
- 4 Reduction to SQL
- 5 Computation of the GMD-Join
- 6 Experimental Evaluation**

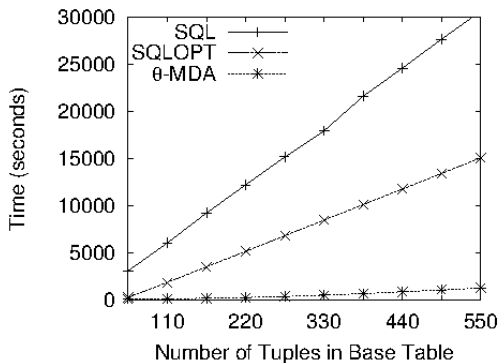
Experimental Evaluation/1

- We use Query Q1 (CntDD, CumCntD, and CumCntDD)
- Varying size of detail table (base table = 550 tuples).



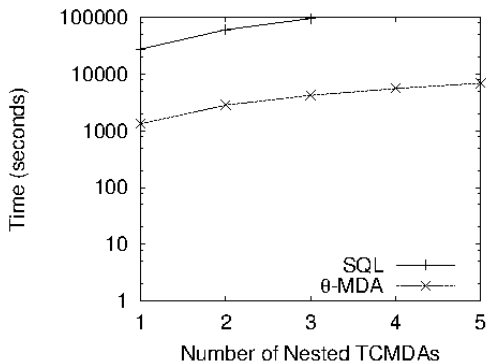
Experimental Evaluation/2

- Varying size of base table (detail table = 50M tuples).



Experimental Evaluation/3

- Varying number of nested GMD-Joins (base table = 550 tuples, detail table = 50M tuples).



Summary

- The GMD-join is an operator for **ad hoc OLAP** queries
 - **Clear separation** between the grouping and the aggregation. This provides flexibility.
 - **Easy formulation** of complex queries.
 - Allows to transform complex predicates into simple predicates.
- **Efficient evaluation** algorithms exist.
 - **Single pass** over detail table.
 - Further optimizations are possible to improve the performance.
- The GMDJ interacts with other algebraic operators (**transformation rules**).
- Intermediate results are typically small (unlike in join-based solutions)
- **Systematic transformation to SQL** is possible, but might be inefficient.
- **2D aggregates** cannot be computed with SQL window functions.