

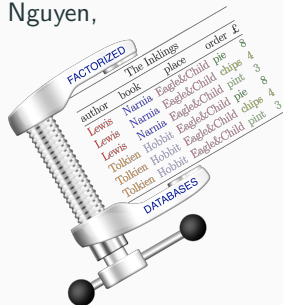
In-Database Factorised Learning

fdbresearch.github.io

Mahmoud Abo Khamis, Hung Ngo, XuanLong Nguyen,
Dan Olteanu, and Maximilian Schleich

December 2017

Logic for Data Science Seminar
Alan Turing Institute



Current Landscape for ML over DB

Factorised Learning over Normalised Data

Learning under Functional Dependencies

General Problem Formulation

Brief Outlook at Current Landscape for ML over DB (1/2)

No integration

- ML & DB distinct tools on the technology stack
- DB exports data as one table, ML imports it in own format
- Spark/PostgreSQL + R supports virtually any ML task
- Most ML over DB solutions operate in this space

Brief Outlook at Current Landscape for ML over DB (1/2)

No integration

- ML & DB distinct tools on the technology stack
- DB exports data as one table, ML imports it in own format
- **Spark/PostgreSQL + R** supports virtually any ML task
- Most ML over DB solutions operate in this space

Loose integration

- Each ML task implemented by a distinct UDF inside DB
- Same running process for DB and ML
- DB computes one table, ML works directly on it
- **MadLib** supports comprehensive library of ML UDFs

Brief Outlook at Current Landscape for ML over DB (2/2)

Unified programming architecture

- One framework for many ML tasks instead of one UDF per task, with possible code reuse across UDFs
- DB computes one table, ML works directly on it
- **Bismark** supports incremental gradient descent for convex programming; up to 100% overhead over specialized UDFs

Brief Outlook at Current Landscape for ML over DB (2/2)

Unified programming architecture

- One framework for many ML tasks instead of one UDF per task, with possible code reuse across UDFs
- DB computes one table, ML works directly on it
- **Bismark** supports incremental gradient descent for convex programming; up to 100% overhead over specialized UDFs

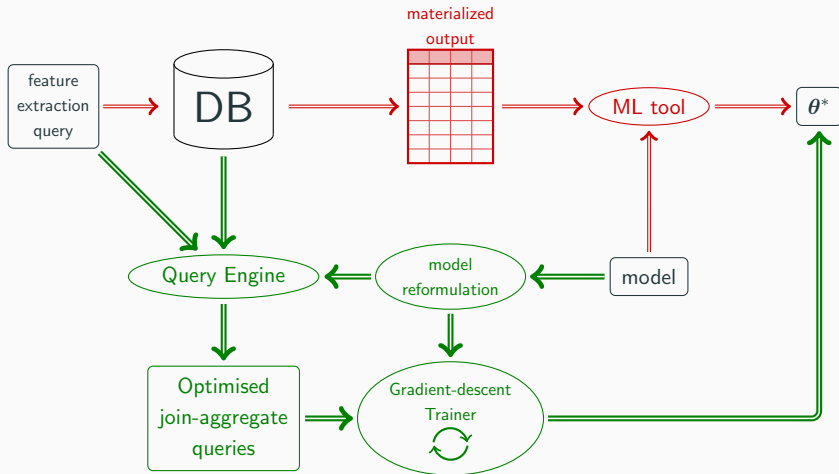
Tight integration $\Rightarrow$ In-Database Analytics

- One evaluation plan for both DB and ML workload; opportunity to push parts of ML tasks past DB joins
- **Morpheus + Hamlet** supports GLM and naïve Bayes
- **Our approach** supports PR/FM, decision trees, ...

In-Database Analytics

- Move the analytics, not the data
 - Avoid expensive data export/import
 - Exploit database technologies
 - Exploit the relational structure (schema, query, dependencies)
 - Build better models using larger datasets and faster
- Cast analytics code as join-aggregate queries
 - Many similar queries that massively share computation
 - Fixpoint computation needed for model convergence

In-database vs. Out-of-database Analytics



Does It Pay Off in Practice?

Retailer dataset (records)		excerpt (17M)	full (86M)
Linear regression			
Features	(cont+categ)	33 + 55	33+3,653
Aggregates	(cont+categ)	595+2,418	595+145k
MadLib	Learn	1,898.35 sec	> 24h
R	Join (PSQL)	50.63 sec	–
	Export/Import	308.83 sec	–
	Learn	490.13 sec	–
Our approach	Join-Aggregate	25.51 sec	380.31 sec
(1core, commodity machine)	Converge (runs)	0.02 (343) sec	8.82 (366) sec
Polynomial regression degree 2			
Features	(cont+categ)	562+2,363	562+141k
Aggregates	(cont+categ)	158k+742k	158k+37M
MadLib	Learn	> 24h	–
Our approach	Join-Aggregate	132.43 sec	1,819.80 sec
(1core, commodity machine)	Converge (runs)	3.27 (321) sec	219.51 (180) sec

Current Landscape for ML over DB

Factorised Learning over Normalised Data

Learning under Functional Dependencies

General Problem Formulation

Unified In-Database Analytics for Optimisation Problems

Our target: retail-planning and forecasting applications

- **Typical databases:** weekly sales, promotions, and products
- **Training dataset:** Result of a feature extraction query
- **Task:** Train model to predict additional demand generated for a product due to promotion
- **Training algorithm:** batch gradient descent
- **ML tasks:** ridge linear regression, polynomial regression, factorisation machines; logistic regression, SVM; PCA.

Typical Retail Example

- Database $I = (R_1, R_2, R_3, R_4, R_5)$
- Feature selection query Q :

$Q(\text{sku}, \text{store}, \text{color}, \text{city}, \text{country}, \text{unitsSold}) \leftarrow$
 $R_1(\text{sku}, \text{store}, \text{day}, \text{unitsSold}), R_2(\text{sku}, \text{color}),$
 $R_3(\text{day}, \text{quarter}), R_4(\text{store}, \text{city}), R_5(\text{city}, \text{country}).$

- Free variables
 - Categorical (qualitative):
 $F = \{\text{sku}, \text{store}, \text{color}, \text{city}, \text{country}\}.$
 - Continuous (quantitative): $\text{unitsSold}.$
- Bound variables
 - Categorical (qualitative): $B = \{\text{day}, \text{quarter}\}$

Typical Retail Example

- We learn the ridge linear regression model

$$\langle \boldsymbol{\theta}, \mathbf{x} \rangle = \sum_{f \in F} \langle \boldsymbol{\theta}_f, \mathbf{x}_f \rangle$$

- Training dataset: $D = Q(I)$
- Feature vector $\mathbf{x}$ and response $y = \text{unitsSold}$
- The parameters $\boldsymbol{\theta}$ obtained by minimising the objective function:

$$J(\boldsymbol{\theta}) = \underbrace{\frac{1}{2|D|} \sum_{(\mathbf{x}, y) \in D} (\langle \boldsymbol{\theta}, \mathbf{x} \rangle - y)^2}_{\text{least square loss}} + \underbrace{\|\boldsymbol{\theta}\|_2^2}_{\ell_2\text{-regulariser}}$$

Side Note: One-hot Encoding of Categorical Variables

- **Continuous** variables are mapped to scalars
 - $y_{unitsSold} \in \mathbb{R}$.
- **Categorical** variables are mapped to indicator vectors
 - country has categories vietnam and england
 - country is then mapped to an indicator vector
$$\mathbf{x}_{country} = [x_{vietnam}, x_{england}]^T \in (\{0, 1\}^2)^T.$$
 - $\mathbf{x}_{country} = [0, 1]^T$ for a tuple with country = ‘‘england’’

This encoding leads to wide training datasets and many 0s

From Optimisation to Sum-Product Queries

We can solve $\theta^* := \arg \min_{\theta} J(\theta)$ by repeatedly updating θ in the direction of the gradient until convergence:

$$\theta := \theta - \alpha \cdot \nabla J(\theta).$$

From Optimisation to Sum-Product Queries

We can solve $\theta^* := \arg \min_{\theta} J(\theta)$ by repeatedly updating θ in the direction of the gradient until convergence:

$$\theta := \theta - \alpha \cdot \nabla J(\theta).$$

Define the matrix $\Sigma = (\sigma_{ij})_{i,j \in [|F|]}$, the vector $\mathbf{c} = (c_i)_{i \in [|F|]}$, and the scalar s_Y :

$$\sigma_{ij} = \frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} \mathbf{x}_i \mathbf{x}_j^{\top} \quad \mathbf{c}_i = \frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} y \cdot \mathbf{x}_i \quad s_Y = \frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} y^2.$$

From Optimisation to Sum-Product Queries

We can solve $\theta^* := \arg \min_{\theta} J(\theta)$ by repeatedly updating θ in the direction of the gradient until convergence:

$$\theta := \theta - \alpha \cdot \nabla J(\theta).$$

Define the matrix $\Sigma = (\sigma_{ij})_{i,j \in [|F|]}$, the vector $\mathbf{c} = (c_i)_{i \in [|F|]}$, and the scalar s_Y :

$$\sigma_{ij} = \frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} \mathbf{x}_i \mathbf{x}_j^{\top} \quad \mathbf{c}_i = \frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} y \cdot \mathbf{x}_i \quad s_Y = \frac{1}{|D|} \sum_{(\mathbf{x}, y) \in D} y^2.$$

Then,

$$\begin{aligned} J(\theta) &= \frac{1}{2|D|} \sum_{(\mathbf{x}, y) \in D} (\langle \theta, \mathbf{x} \rangle - y)^2 + \frac{\lambda}{2} \|\theta\|_2^2 \\ &= \frac{1}{2} \theta^{\top} \Sigma \theta - \langle \theta, \mathbf{c} \rangle + \frac{s_Y}{2} + \frac{\lambda}{2} \|\theta\|_2^2 \end{aligned}$$

Σ , c , s_Y can be Expressed as SQL Queries

SQL queries for $\sigma_{ij} = \frac{1}{|D|} \sum_{(x,y) \in D} \mathbf{x}_i \mathbf{x}_j^\top$ (w/o factor $\frac{1}{|D|}$):

Σ , c , s_Y can be Expressed as SQL Queries

SQL queries for $\sigma_{ij} = \frac{1}{|D|} \sum_{(x,y) \in D} \mathbf{x}_i \mathbf{x}_j^\top$ (w/o factor $\frac{1}{|D|}$):

- x_i, x_j continuous $\Rightarrow$ no group-by variable

```
SELECT SUM ( $x_i * x_j$ ) FROM  $D$ ;
```

where D is the natural join of tables R_1 to R_5 in our example.

Σ , c , s_Y can be Expressed as SQL Queries

SQL queries for $\sigma_{ij} = \frac{1}{|D|} \sum_{(x,y) \in D} \mathbf{x}_i \mathbf{x}_j^\top$ (w/o factor $\frac{1}{|D|}$):

- x_i, x_j continuous $\Rightarrow$ no group-by variable

```
SELECT SUM ( $x_i * x_j$ ) FROM  $D$ ;
```

- x_i categorical, x_j continuous $\Rightarrow$ one group-by variable

```
SELECT  $x_i$ , SUM( $x_j$ ) FROM  $D$  GROUP BY  $x_i$ ;
```

where D is the natural join of tables R_1 to R_5 in our example.

Σ , c , s_Y can be Expressed as SQL Queries

SQL queries for $\sigma_{ij} = \frac{1}{|D|} \sum_{(x,y) \in D} \mathbf{x}_i \mathbf{x}_j^\top$ (w/o factor $\frac{1}{|D|}$):

- x_i, x_j continuous $\Rightarrow$ no group-by variable

```
SELECT SUM ( $x_i * x_j$ ) FROM  $D$ ;
```

- x_i categorical, x_j continuous $\Rightarrow$ one group-by variable

```
SELECT  $x_i$ , SUM( $x_j$ ) FROM  $D$  GROUP BY  $x_i$ ;
```

- x_i, x_j categorical $\Rightarrow$ two group-by variables

```
SELECT  $x_i, x_j$ , SUM(1) FROM  $D$  GROUP BY  $x_i, x_j$ ;
```

where D is the natural join of tables R_1 to R_5 in our example.

Σ , c , s_Y can be Expressed as SQL Queries

SQL queries for $\sigma_{ij} = \frac{1}{|D|} \sum_{(x,y) \in D} \mathbf{x}_i \mathbf{x}_j^\top$ (w/o factor $\frac{1}{|D|}$):

- x_i, x_j continuous $\Rightarrow$ no group-by variable

```
SELECT SUM ( $x_i * x_j$ ) FROM  $D$ ;
```

- x_i categorical, x_j continuous $\Rightarrow$ one group-by variable

```
SELECT  $x_i$ , SUM( $x_j$ ) FROM  $D$  GROUP BY  $x_i$ ;
```

- x_i, x_j categorical $\Rightarrow$ two group-by variables

```
SELECT  $x_i, x_j$ , SUM(1) FROM  $D$  GROUP BY  $x_i, x_j$ ;
```

where D is the natural join of tables R_1 to R_5 in our example.

This query encoding avoids drawbacks of one-hot encoding

How To Compute Efficiently These Join-Aggregate Queries?

Factorised Query Computation by Example

Orders (O for short)			Dish (D for short)		Items (I for short)	
customer	day	dish	dish	item	item	price
Elise	Monday	burger	burger	patty	patty	6
Elise	Friday	burger	burger	onion	onion	2
Steve	Friday	hotdog	burger	bun	bun	2
Joe	Friday	hotdog	hotdog	bun	sausage	4
			hotdog	onion		
			hotdog	sausage		

Factorised Query Computation by Example

Orders (O for short)			Dish (D for short)		Items (I for short)	
customer	day	dish	dish	item	item	price
Elise	Monday	burger	burger	patty	patty	6
Elise	Friday	burger	burger	onion	onion	2
Steve	Friday	hotdog	burger	bun	bun	2
Joe	Friday	hotdog	hotdog	bun	sausage	4
			hotdog	onion		
			hotdog	sausage		

Consider the natural join of the above relations:

O(customer, day, **dish**), D(**dish**, **item**), I(**item**, price)

customer	day	dish	item	price
Elise	Monday	burger	patty	6
Elise	Monday	burger	onion	2
Elise	Monday	burger	bun	2
Elise	Friday	burger	patty	6
Elise	Friday	burger	onion	2
Elise	Friday	burger	bun	2
...	...	...	...	...

Factorised Query Computation by Example

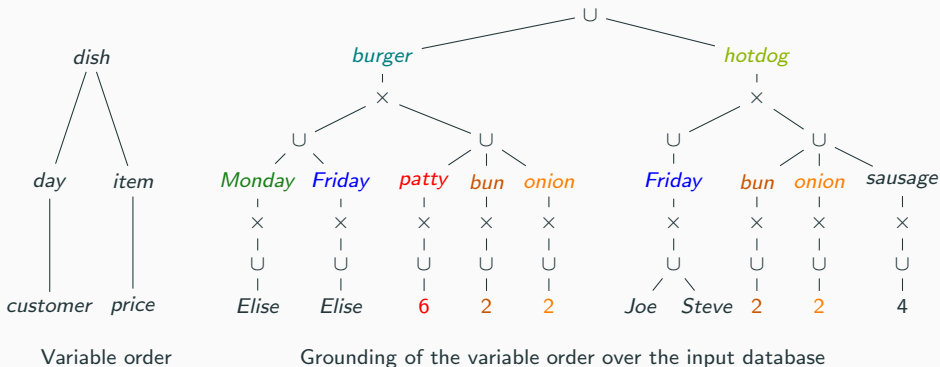
$$O(\text{customer, day, dish}), D(\text{dish, item}), I(\text{item, price})$$

customer	day	dish	item	price
Elise	Monday	burger	patty	6
Elise	Monday	burger	onion	2
Elise	Monday	burger	bun	2
Elise	Friday	burger	patty	6
Elise	Friday	burger	onion	2
Elise	Friday	burger	bun	2
...	...	...	...	...

An algebraic encoding uses product ($\times$), union ($\cup$), and values:

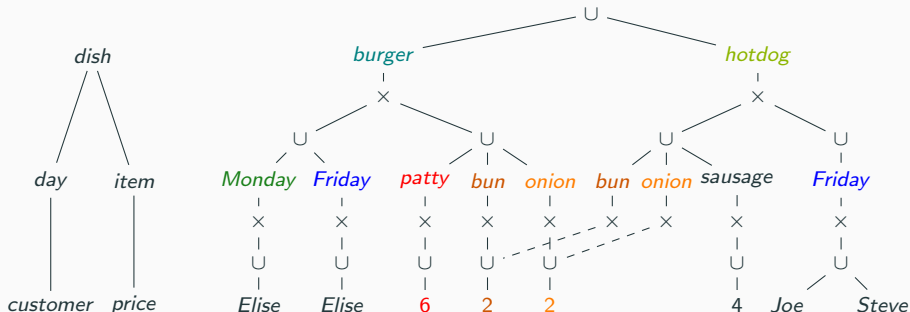
Elise $\times$ *Monday* $\times$ *burger* $\times$ *patty* $\times$ *6* $\cup$
Elise $\times$ *Monday* $\times$ *burger* $\times$ *onion* $\times$ *2* $\cup$
Elise $\times$ *Monday* $\times$ *burger* $\times$ *bun* $\times$ *2* $\cup$
Elise $\times$ *Friday* $\times$ *burger* $\times$ *patty* $\times$ *6* $\cup$
Elise $\times$ *Friday* $\times$ *burger* $\times$ *onion* $\times$ *2* $\cup$
Elise $\times$ *Friday* $\times$ *burger* $\times$ *bun* $\times$ *2* $\cup \dots$

Factorised Join



There are several **algebraically equivalent** factorised joins defined by distributivity of product over union and their commutativity.

.. Now with Further Compression

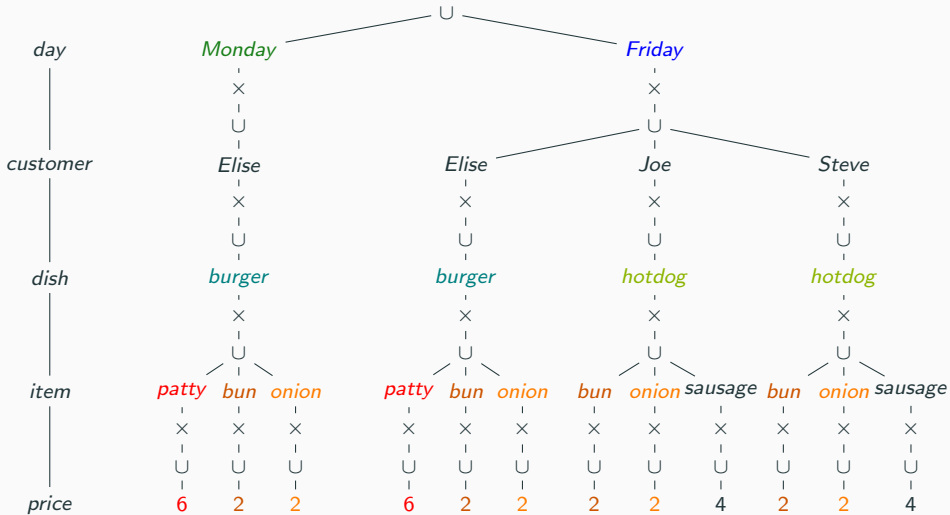


Observation:

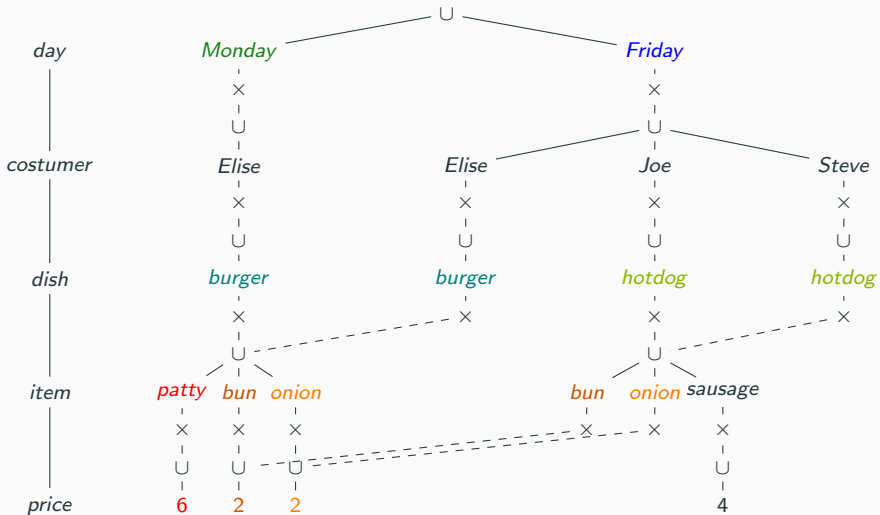
- price is under item, which is under dish, but only *depends* on item,
- .. so the same price appears under an item *regardless* of the dish.

Idea: *Cache* price for a specific item and avoid repetition!

Same Data, Different Factorisation

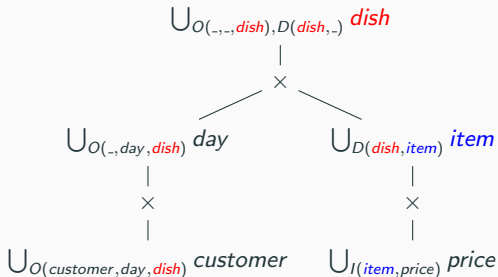


.. and Further Compressed



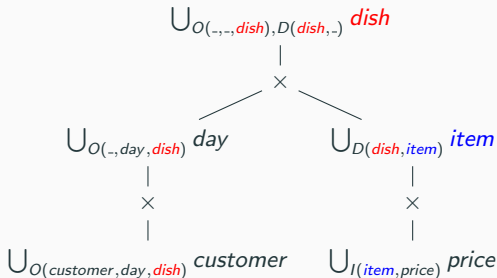
Grounding Variable Orders to Factorised Joins

Our join: $O(\text{customer}, \text{day}, \text{dish})$, $D(\text{dish}, \text{item})$, $I(\text{item}, \text{price})$
can be grounded to a factorised join as follows:



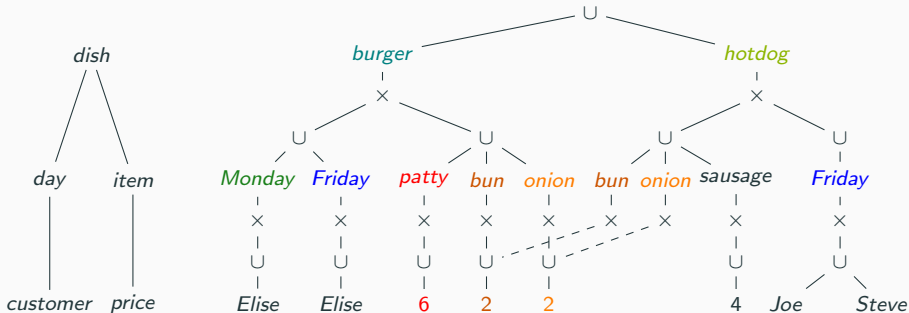
This grounding follows the previous variable order.

Grounding Variable Orders to Factorised Joins



- Relations sorted following topological order of the variable order
- Intersection of O and D on *dish* in time $\tilde{O}(\min(|\pi_{\text{dish}} O|, |\pi_{\text{dish}} D|))$
- The remaining operations are lookups in the relations, where we first fix the *dish* value and then the *day* and *item* values

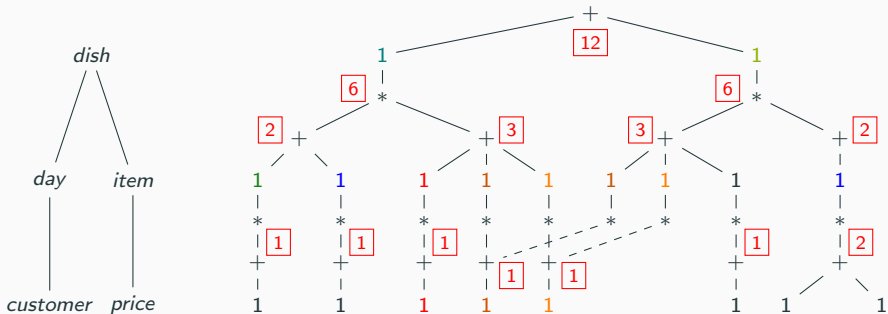
Factorising the Computation of Aggregates (1/2)



COUNT(*) computed in one pass over the factorisation:

- values $\mapsto 1$,
- $\cup \mapsto +$,
- $\times \mapsto *$.

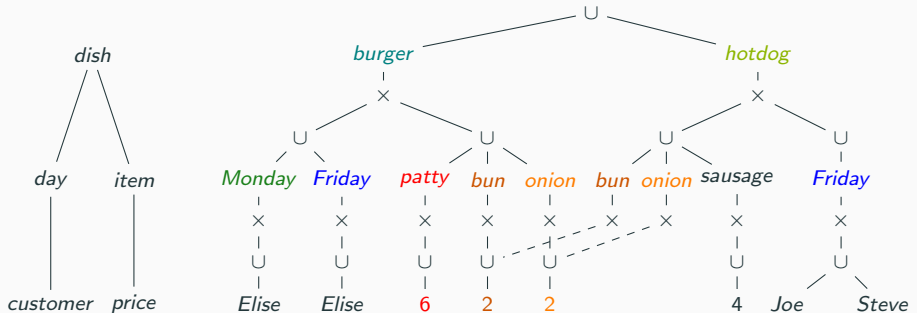
Factorising the Computation of Aggregates (1/2)



COUNT(*) computed in one pass over the factorisation:

- $\text{values} \mapsto 1$,
- $\cup \mapsto +$,
- $\times \mapsto *$.

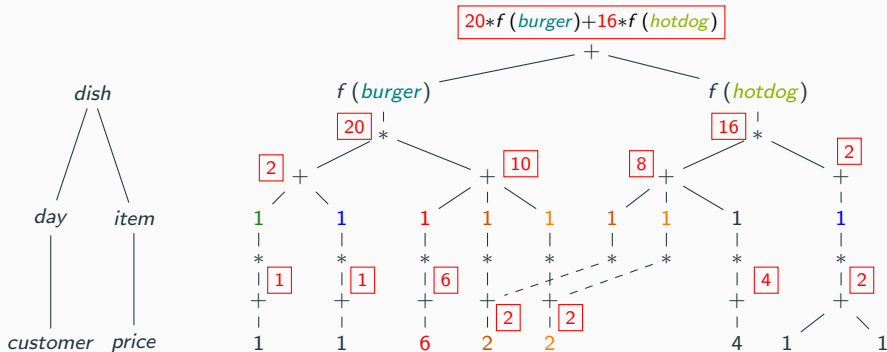
Factorising the Computation of Aggregates (2/2)



$\text{SUM}(\text{dish} * \text{price})$ computed in one pass over the factorisation:

- Assume there is a function f that turns dish into reals.
- All values except for dish & price $\mapsto 1$,
- $U \mapsto +$,
- $x \mapsto *$.

Factorising the Computation of Aggregates (2/2)



SUM(dish * price) computed in one pass over the factorisation:

- Assume there is a function f that turns dish into reals.
- All values except for dish & price $\mapsto 1$,
- $\cup \mapsto +$,
- $\times \mapsto *$.

Current Landscape for ML over DB

Factorised Learning over Normalised Data

Learning under Functional Dependencies

General Problem Formulation

Model Reparameterisation using Functional Dependencies

Consider the functional dependency $\text{city} \rightarrow \text{country}$ and

- country categories: vietnam, england
- city categories: saigon, hanoi, oxford, leeds, bristol

The one-hot encoding enforces the following identities:

- $x_{\text{vietnam}} = x_{\text{saigon}} + x_{\text{hanoi}}$

country is vietnam $\Rightarrow$ city is either saigon or hanoi

$$x_{\text{vietnam}} = 1 \Rightarrow \text{either } x_{\text{saigon}} = 1 \text{ or } x_{\text{hanoi}} = 1$$

- $x_{\text{england}} = x_{\text{oxford}} + x_{\text{leeds}} + x_{\text{bristol}}$

country is england $\Rightarrow$ city is either oxford, leeds, or bristol

$$x_{\text{england}} = 1 \Rightarrow \text{either } x_{\text{oxford}} = 1 \text{ or } x_{\text{leeds}} = 1 \text{ or } x_{\text{bristol}} = 1$$

Model Reparameterisation using Functional Dependencies

- Identities due to one-hot encoding

$$x_{\text{vietnam}} = x_{\text{saigon}} + x_{\text{hanoi}}$$

$$x_{\text{england}} = x_{\text{oxford}} + x_{\text{leeds}} + x_{\text{bristol}}$$

- Encode $\mathbf{x}_{\text{country}}$ as $\mathbf{x}_{\text{country}} = \mathbf{R}\mathbf{x}_{\text{city}}$, where

	saigon	hanoi	oxford	leeds	bristol	
$\mathbf{R} =$	1	1	0	0	0	vietnam
	0	0	1	1	1	england

For instance, if city is saigon, i.e., $\mathbf{x}_{\text{city}} = [1, 0, 0, 0, 0]^T$, then country is vietnam, i.e., $\mathbf{x}_{\text{country}} = \mathbf{R}\mathbf{x}_{\text{city}} = [1, 0]^T$.

$$\begin{bmatrix} 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$$

Model Reparameterisation using Functional Dependencies

- Functional dependency: $\text{city} \rightarrow \text{country}$
- $\mathbf{x}_{\text{country}} = \mathbf{R}\mathbf{x}_{\text{city}}$
- Replace all occurrences of $\mathbf{x}_{\text{country}}$ by $\mathbf{R}\mathbf{x}_{\text{city}}$:

$$\begin{aligned} & \sum_{f \in F - \{\text{city}, \text{country}\}} \langle \boldsymbol{\theta}_f, \mathbf{x}_f \rangle + \langle \boldsymbol{\theta}_{\text{country}}, \mathbf{x}_{\text{country}} \rangle + \langle \boldsymbol{\theta}_{\text{city}}, \mathbf{x}_{\text{city}} \rangle \\ = & \sum_{f \in F - \{\text{city}, \text{country}\}} \langle \boldsymbol{\theta}_f, \mathbf{x}_f \rangle + \langle \boldsymbol{\theta}_{\text{country}}, \mathbf{R}\mathbf{x}_{\text{city}} \rangle + \langle \boldsymbol{\theta}_{\text{city}}, \mathbf{x}_{\text{city}} \rangle \\ = & \sum_{f \in F - \{\text{city}, \text{country}\}} \langle \boldsymbol{\theta}_f, \mathbf{x}_f \rangle + \left\langle \underbrace{\mathbf{R}^\top \boldsymbol{\theta}_{\text{country}} + \boldsymbol{\theta}_{\text{city}}}_{\gamma_{\text{city}}}, \mathbf{x}_{\text{city}} \right\rangle \end{aligned}$$

Model Reparameterisation using Functional Dependencies

- Functional dependency: $\text{city} \rightarrow \text{country}$
- $\mathbf{x}_{\text{country}} = \mathbf{R}\mathbf{x}_{\text{city}}$
- Replace all occurrences of $\mathbf{x}_{\text{country}}$ by $\mathbf{R}\mathbf{x}_{\text{city}}$:

$$\begin{aligned} & \sum_{f \in F - \{\text{city}, \text{country}\}} \langle \boldsymbol{\theta}_f, \mathbf{x}_f \rangle + \langle \boldsymbol{\theta}_{\text{country}}, \mathbf{x}_{\text{country}} \rangle + \langle \boldsymbol{\theta}_{\text{city}}, \mathbf{x}_{\text{city}} \rangle \\ &= \sum_{f \in F - \{\text{city}, \text{country}\}} \langle \boldsymbol{\theta}_f, \mathbf{x}_f \rangle + \langle \boldsymbol{\theta}_{\text{country}}, \mathbf{R}\mathbf{x}_{\text{city}} \rangle + \langle \boldsymbol{\theta}_{\text{city}}, \mathbf{x}_{\text{city}} \rangle \\ &= \sum_{f \in F - \{\text{city}, \text{country}\}} \langle \boldsymbol{\theta}_f, \mathbf{x}_f \rangle + \left\langle \underbrace{\mathbf{R}^\top \boldsymbol{\theta}_{\text{country}} + \boldsymbol{\theta}_{\text{city}}}_{\boldsymbol{\gamma}_{\text{city}}}, \mathbf{x}_{\text{city}} \right\rangle \end{aligned}$$

- We avoid the computation of the aggregates over $\mathbf{x}_{\text{country}}$.
- We reparameterise and ignore parameters $\boldsymbol{\theta}_{\text{country}}$.
- What about the penalty term in the loss function?

Model Reparameterisation using Functional Dependencies

- Functional dependency: $\text{city} \rightarrow \text{country}$
- $\mathbf{x}_{\text{country}} = \mathbf{R}\mathbf{x}_{\text{city}} \quad \gamma_{\text{city}} = \mathbf{R}^\top \boldsymbol{\theta}_{\text{country}} + \boldsymbol{\theta}_{\text{city}}$

- Rewrite the penalty term

$$\|\boldsymbol{\theta}\|_2^2 = \sum_{j \neq \text{city}} \|\boldsymbol{\theta}_j\|_2^2 + \left\| \gamma_{\text{city}} - \mathbf{R}^\top \boldsymbol{\theta}_{\text{country}} \right\|_2^2 + \|\boldsymbol{\theta}_{\text{country}}\|_2^2$$

- Optimise out $\boldsymbol{\theta}_{\text{country}}$ by expressing it in terms of γ_{city} :

$$\boldsymbol{\theta}_{\text{country}} = (\mathbf{I}_{\text{country}} + \mathbf{R}\mathbf{R}^\top)^{-1} \mathbf{R} \gamma_{\text{city}} = \mathbf{R}(\mathbf{I}_{\text{city}} + \mathbf{R}^\top \mathbf{R})^{-1} \gamma_{\text{city}}$$

- The penalty term becomes

$$\|\boldsymbol{\theta}\|_2^2 = \sum_{j \neq \text{city}} \|\boldsymbol{\theta}_j\|_2^2 + \left\langle (\mathbf{I}_{\text{city}} + \mathbf{R}^\top \mathbf{R})^{-1} \gamma_{\text{city}}, \gamma_{\text{city}} \right\rangle$$

Current Landscape for ML over DB

Factorised Learning over Normalised Data

Learning under Functional Dependencies

General Problem Formulation

General Problem Formulation

We want to solve $\theta^* := \arg \min_{\theta} J(\theta)$, where

$$J(\theta) := \sum_{(\mathbf{x}, y) \in D} \mathcal{L}(\langle g(\theta), h(\mathbf{x}) \rangle, y) + \Omega(\theta).$$

- $\theta = (\theta_1, \dots, \theta_p) \in \mathbf{R}^p$ are parameters
- functions $g : \mathbf{R}^p \rightarrow \mathbf{R}^m$ and $h : \mathbf{R}^n \rightarrow \mathbf{R}^m$
 - $g = (g_j)_{j \in [m]}$ is a vector of multivariate polynomials
 - $h = (h_j)_{j \in [m]}$ is a vector of multivariate monomials
- $\mathcal{L}$ is a loss function, Ω is the regulariser
- D is the training dataset with features $\mathbf{x}$ and response y .

Problems: ridge linear regression, polynomial regression,
Factorisation machines; logistic regression, SVM; PCA.

Special Case: Ridge Linear Regression

Under

- square loss $\mathcal{L}$, ℓ_2 -regularisation,
- data points $\mathbf{x} = (x_0, x_1, \dots, x_n, y)$,
- $p = n + 1$ parameters $\boldsymbol{\theta} = (\theta_0, \dots, \theta_n)$,
- $x_0 = 1$ corresponds to the bias parameter θ_0 ,
- identity functions g and h ,

we obtain the following formulation for ridge linear regression:

$$J(\boldsymbol{\theta}) := \frac{1}{2|D|} \sum_{(\mathbf{x}, y) \in D} (\langle \boldsymbol{\theta}, \mathbf{x} \rangle - y)^2 + \frac{\lambda}{2} \|\boldsymbol{\theta}\|_2^2.$$

Special Case: Degree- d Polynomial Regression

Under

- square loss $\mathcal{L}$, ℓ_2 -regularisation,
- data points $\mathbf{x} = (x_0, x_1, \dots, x_n, y)$,
- $p = m = 1 + n + n^2 + \dots + n^d$ parameters $\boldsymbol{\theta} = (\theta_{\mathbf{a}})$, where $\mathbf{a} = (a_1, \dots, a_n)$ with non-negative integers s.t. $\|\mathbf{a}\|_1 \leq d$.
- the components of h are given by $h_{\mathbf{a}}(\mathbf{x}) = \prod_{i=1}^n x_i^{a_i}$,
- $g(\boldsymbol{\theta}) = \boldsymbol{\theta}$,

we obtain the following formulation for polynomial regression:

$$J(\boldsymbol{\theta}) := \frac{1}{2|D|} \sum_{(\mathbf{x}, y) \in D} (\langle g(\boldsymbol{\theta}), h(\mathbf{x}) \rangle - y)^2 + \frac{\lambda}{2} \|\boldsymbol{\theta}\|_2^2.$$

Special Case: Factorisation Machines

Under

- square loss $\mathcal{L}$, ℓ_2 -regularisation,
- data points $\mathbf{x} = (x_0, x_1, \dots, x_n, y)$,
- $p = 1 + n + r \cdot n$ parameters,
- $m = 1 + n + \binom{n}{2}$ features,

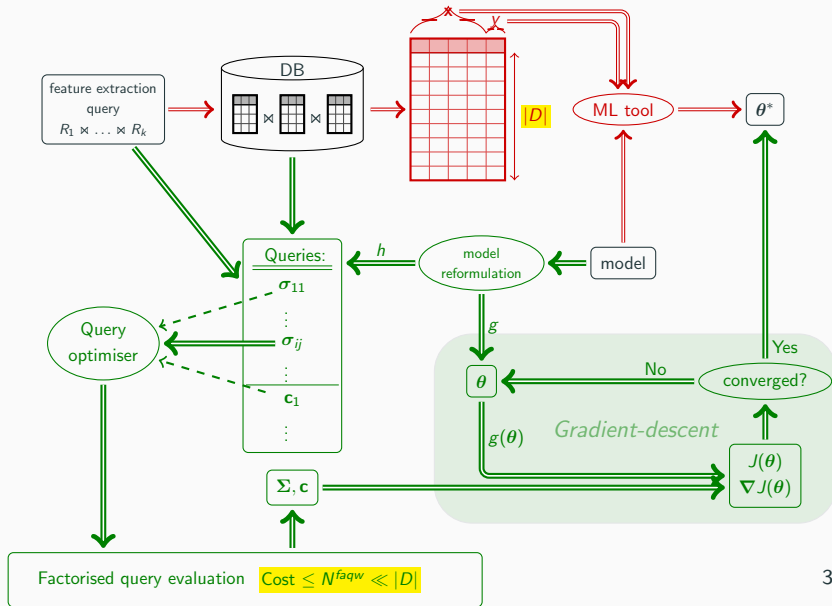
we obtain the following formulation for degree-2 rank- r Factorisation machines:

$$J(\boldsymbol{\theta}) := \frac{1}{2|D|} \sum_{(\mathbf{x}, y) \in D} \left(\sum_{i=0}^n \theta_i x_i + \sum_{\substack{\{i, j\} \in \binom{[n]}{2} \\ \ell \in [r]}} \theta_i^{(\ell)} \theta_j^{(\ell)} x_i x_j - y \right)^2 + \frac{\lambda}{2} \|\boldsymbol{\theta}\|_2^2.$$

Special Case: Classifiers

- Typically, the regulariser is $\frac{\lambda}{2} \|\boldsymbol{\theta}\|_2^2$
- The response is binary: $y \in \{\pm 1\}$
- The loss function $\mathcal{L}(\gamma, y)$, where $\gamma := \langle g(\boldsymbol{\theta}), h(\mathbf{x}) \rangle$ is
 - $\mathcal{L}(\gamma, y) = \max\{1 - y\gamma, 0\}$ for support vector machines,
 - $\mathcal{L}(\gamma, y) = \log(1 + e^{-y\gamma})$ for logistic regression,
 - $\mathcal{L}(\gamma, y) = e^{-y\gamma}$ for Adaboost.

Zoom-in: In-database vs. Out-of-database Learning



Thank you!