

HI!

CS 744: BISMARCK

Shivaram Venkataraman

Fall 2019

ADMINISTRIVIA

- Assignment 2 out!
- Project groups extension → till Oct 7th
- OH / Setup meeting by email Oct 17th Introduction

COURSE PROJECT PROPOSAL

Propose topic, group (2 sentences) – Oct 7

Project Proposal (2 pages) – Oct 17

Introduction

Related Work

Timeline (with eval plan)

WRITING AN INTRODUCTION

1-2 paras: what is the problem you are solving
why is it important (need citations) } → *Crucial*

1-2 paras: How other people solve and why they fall short
↳ *Target audience: Familiar with literature*

1-2 paras: How do you plan on solving it and why your approach is better

1 para: Anticipated results or what experiments you will use

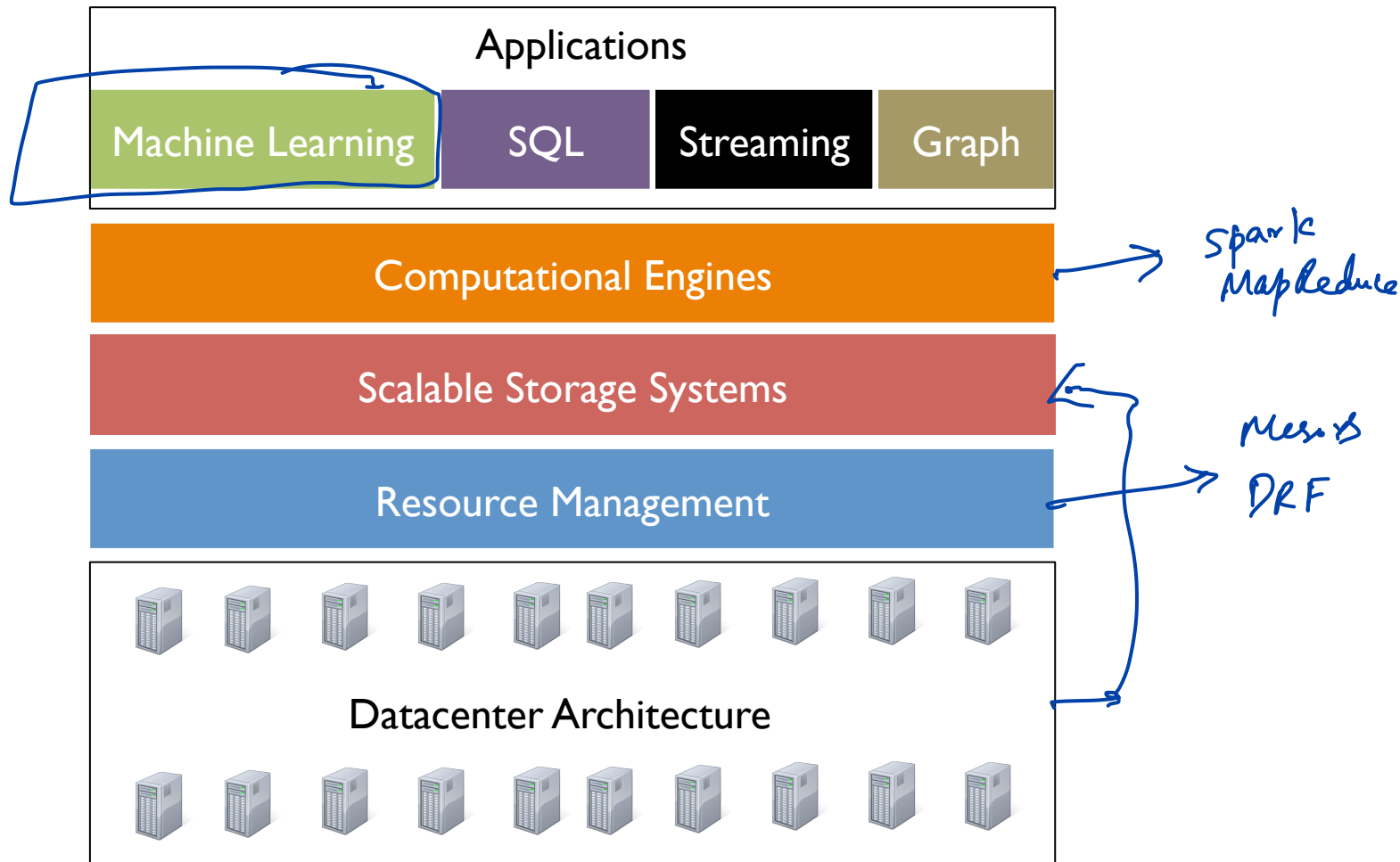
WRITING RELATED WORK

Group related work into two/three buckets
(1-2 para per bucket)

Scheduling *Frameworks*

Explain what the papers / projects do
Why are they different / insufficient

ML Frameworks:



MACHINE LEARNING

Classification



IMAGENET



Recommendation



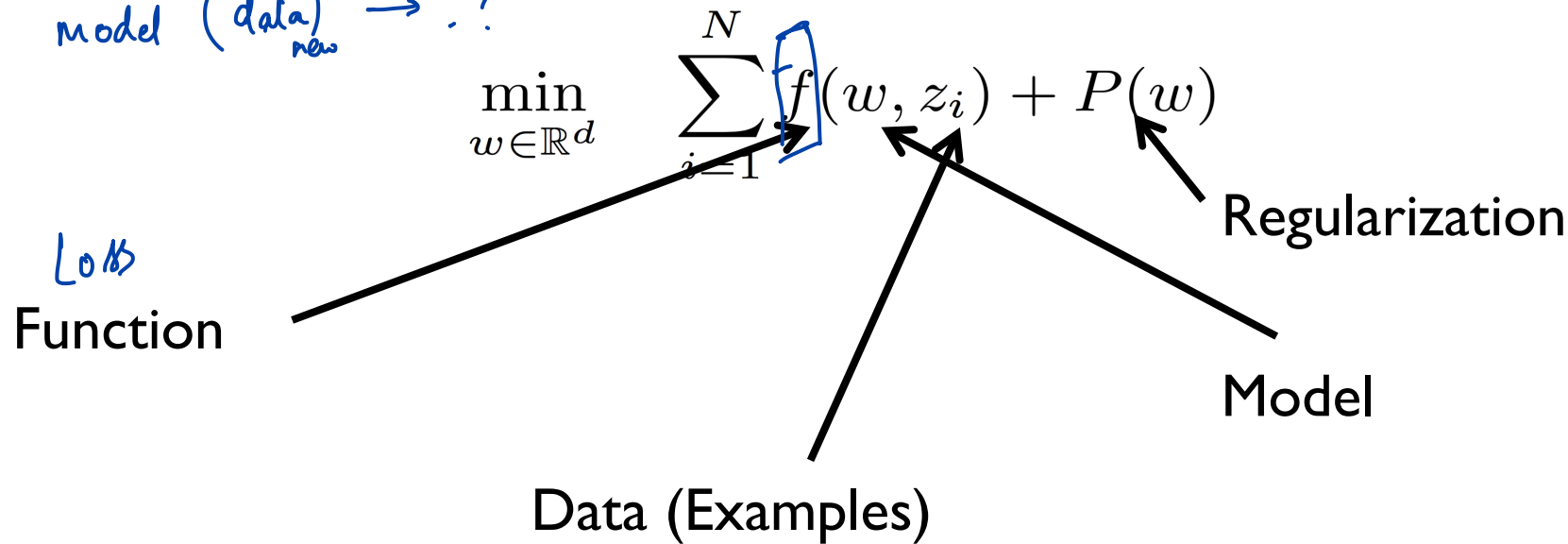
amazon

Challenge



Supervised learning
Training data $\rightarrow$ labels
model (data)_{new} $\rightarrow$??

OPTIMIZATION



CONVEX OPTIMIZATION

$$\min_{w \in \mathbb{R}^d} \sum_{i=1}^N f(w, z_i) + P(w)$$


What is convex ?

Linear Regression, Linear SVM

Kernel SVMs, Logistic Regression,

What is not convex ?

Graph mining, Deep Learning

GRADIENT DESCENT

$$\underbrace{w^{(k+1)}}_{\text{random initialization}} = \underbrace{w^{(k)}}_{\text{model at the } k^{\text{th}} \text{ iteration}} - \underbrace{\alpha_k}_{\text{step size}} \nabla f(w^{(k)})$$

Initialize w → random initialization

For many iterations:

Compute Gradient → How many data points?

Update model

End

GD: All of your training data (n)

IGD: One training data point (1)

SGD: Stochastic Gradient (B ≈ 128 or 512)

INCREMENTAL GRADIENT DESCENT

$$w^{(k+1)} = w^{(k)} - \alpha_k \nabla f_{\eta(k)}(w^{(k)})$$

Initialize w

For many iterations:

Pick one point

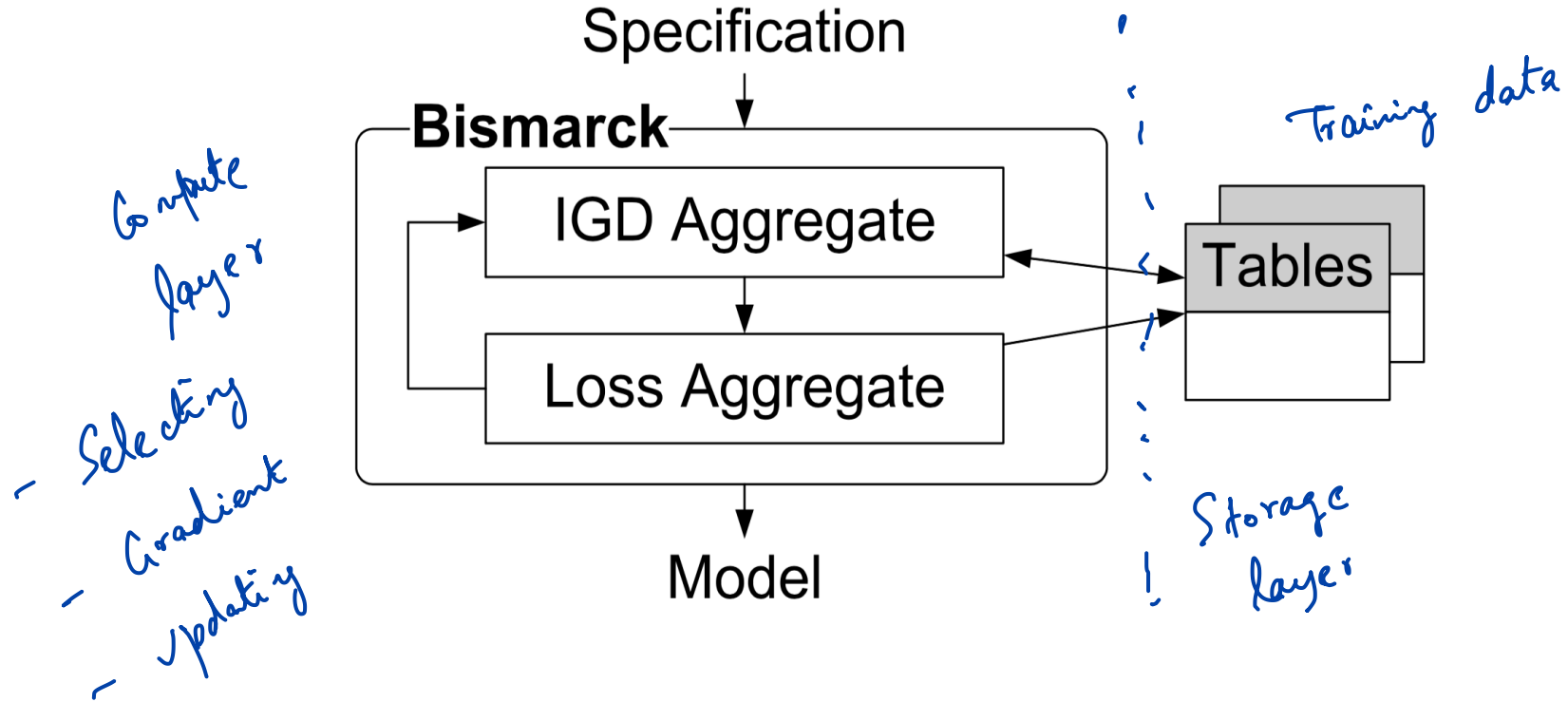
Compute Gradient

Update model

End

Analytics Task	Objective
Logistic Regression (LR)	$\sum_i \log(1 + \exp(-y_i w^T x_i)) + \mu \ \vec{w}\ _1$
Classification (SVM)	$\sum_i (1 - y_i w^T x_i)_+ + \mu \ \vec{w}\ _1$
Recommendation (LMF)	$\sum_{(i,j) \in \Omega} (L_i^T R_j - M_{ij})^2 + \mu \ L, R\ _F^2$
Labeling (CRF) [48]	$\sum_k \left[\sum_j w_j F_j(y_k, x_k) - \log Z(x_k) \right]$
Kalman Filters	$\sum_{t=1}^T \ Cw_t - f(y_t)\ _2^2 + \ w_t - Aw_{t-1}\ _2^2$
Portfolio Optimization	$p^T w + w^T \Sigma w \quad \text{s.t.} \quad w \in \Delta$

BISMARCK ARCHITECTURE



BISMARCK: USER DEFINED AGGREGATE

Three steps:

1. initialize(state) $\hat{=}$ initialize the model
 $\hookrightarrow$ zero weights / random weights

2. transition(state, data)
 $\xrightarrow{\epsilon} \langle \text{data}, \text{label} \rangle$
 $\xrightarrow{\quad} \text{model}$
update model based on data

3. terminate(state)
 $\hookrightarrow$ Loss value $< \|\epsilon\|$
 $\hookrightarrow$ Fixed num iterations / epochs

BISMARCK: LOGISTIC REGRESSION

```
LR_Transition(ModelCoef *w, Example e) { ...  
    wx = Dot_Product(w, e.x);  
    sig = Sigmoid(-wx * e.y);  
    c = stepsize * e.y * sig;  
    Scale_And_Add(w, e.x, c); ... }
```

$[1 | 2 | \dots | n | \dots | 1 | \dots | 1]$

DATA ORDERING

$$n(t_{\text{shuffle}} + t_{\text{epoch}}) \leq \sum_{i=1}^n t_{\text{epoch}} + t_{\text{shuffle}}$$

Random sampling

- Sample **without replacement**
- Shuffle the data after each **epoch**

Shuffle once

- Avoids pathological ordering
- Much cheaper

while iter:

Pick data pt

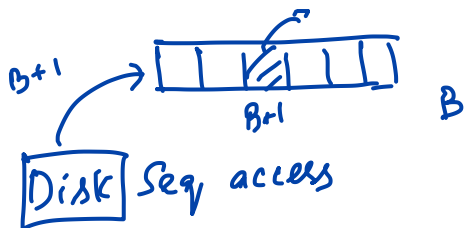
Grad
Update

end

$[1 | \dots | 1 | \dots | 100]$
↓ shuffle

$[51 | 42 | 33 | \dots | 17]$

$[88 | 22 | \dots | 13]$
epoch



first B

RESERVOIR SAMPLING

Select first m items

On the k^{th} additional item

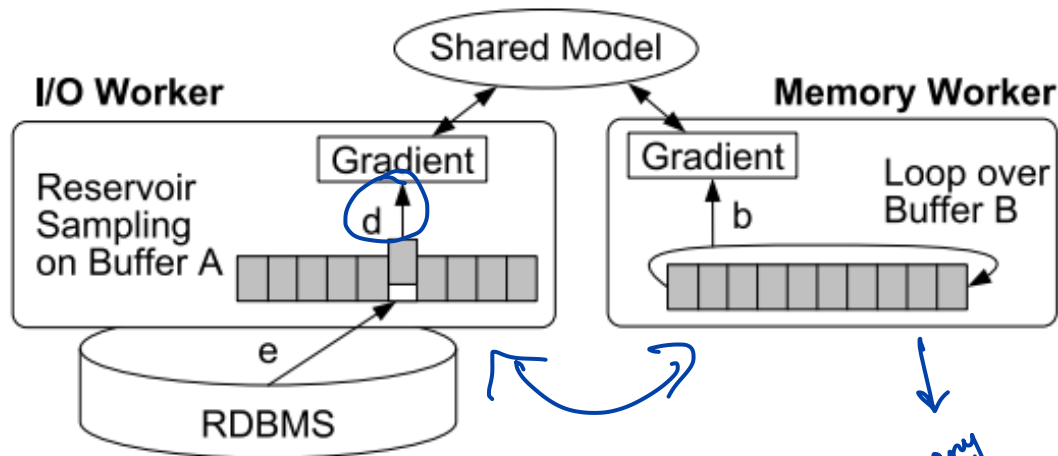
$s = \text{random in } [0, m + k)$

if $s < m$

Put in slot s

else

Drop the item



Reservoir sampling eliminated ICD
→ Keys

many iterations over fixed buffer

Async update
 $T_1: w_{k+1}$
 w_{k+2}

PARALLEL GRADIENTS

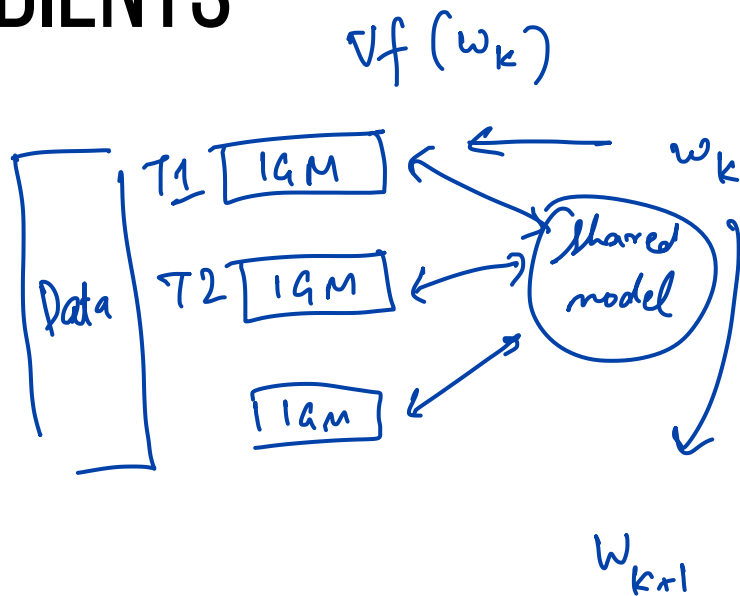
Shared Memory:

- Compute gradients in parallel
- Average their updates
- Or update in parallel
 - Locks?

- AIG: Atomic Increment Instructions

↳ fast, implemented in hardware

- Lock-free: Iteration times very fast, conflicts?



DISCUSSION

<https://forms.gle/nFNEi2NZMNhZioIf7>

How would an implementation of GD look in Spark? Try to sketch an implementation. What would be similar / different to Bismarck?

```
data: sc.textFile("-")
data.shuffle()
→ grad: data.map(x => gradComp(x))
           .reduce(_+_)
```

model

shuffle everytime?

```
model := model - step * grad
→ bop
```

broadcast

check for termination

shared memory
↳ no broadcast
↳ AIC2 & lock free

What are some ML scenarios where Bismarck architecture might prove to be limited?

Non Convex : Converge ? shuffle once
lock free

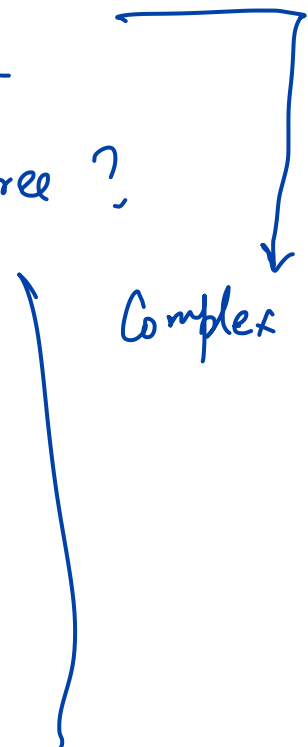
skewed dataset : Conflicts $\rightarrow$ lock-free ?

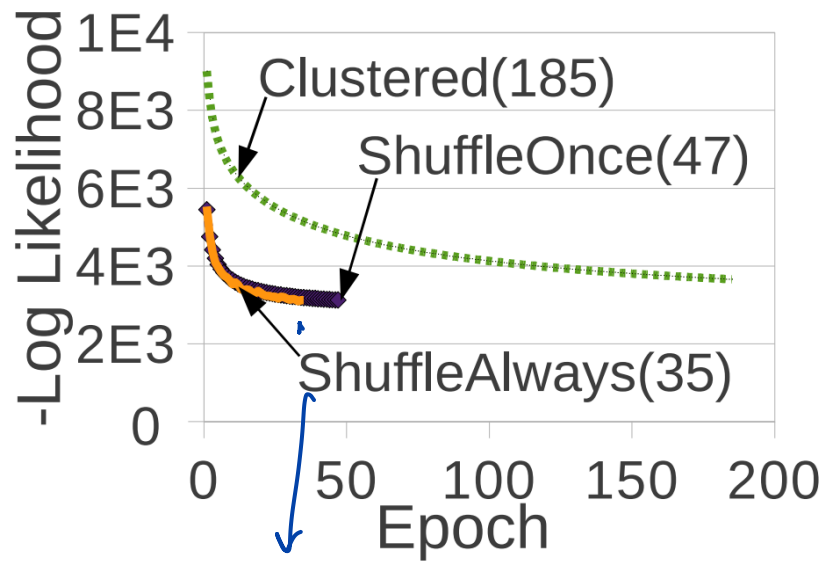
Overfitting ?

Model fits in memory :

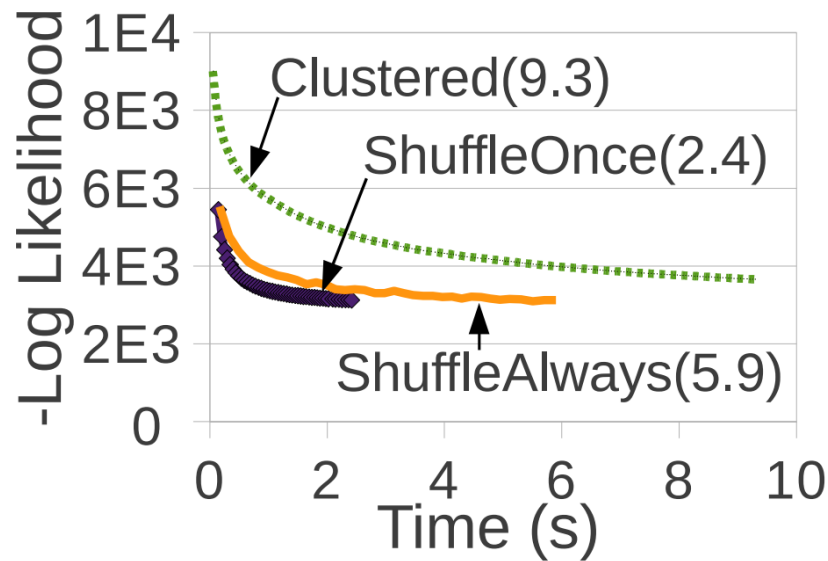
11th ism limited ?

Complex model





leak
epoch



NEXT STEPS

Next class: Parameter Server

Assignment 2 out!

Project Proposal

Groups by Oct 7

2 pager by Oct 17