

Good morning!

CS 744: PARAMETER SERVERS

Shivaram Venkataraman

Fall 2019

ADMINISTRIVIA

- Assignment 2 is out!
- Course project groups due Oct 7
- Introductions due Oct 17

Applications

Machine Learning

SQL

Streaming

Graph

Computational Engines

Scalable Storage Systems

Resource Management



Datacenter Architecture



Bismarck

Supervised learning ←
Unified Interface

Shared memory
Model fits in memory

IGD { Sampling / ordering
Parallelism

Machine Learning

MOTIVATION

- Large training data 1TB to 1PB
- Models with 10^9 to 10^{12} parameters
8GB 8TB if 8 bytes per parameter
- Goals
 - Efficient communication
 - Flexible synchronization
 - Elastic Scalability
 - Fault Tolerance and Durability

EXAMPLE WORKLOAD

$$w_{t+1} = w_t - \gamma \sum_{i=1}^k g_i$$

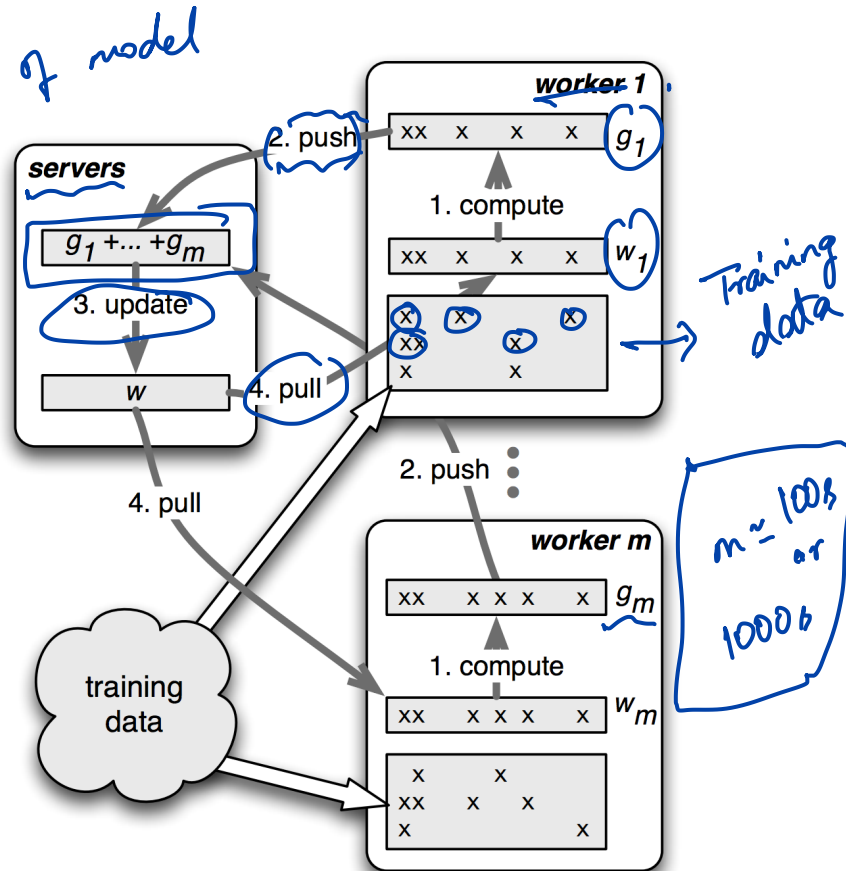
Ad Click Prediction

- Trillions of clicks per day
- Very sparse feature vectors

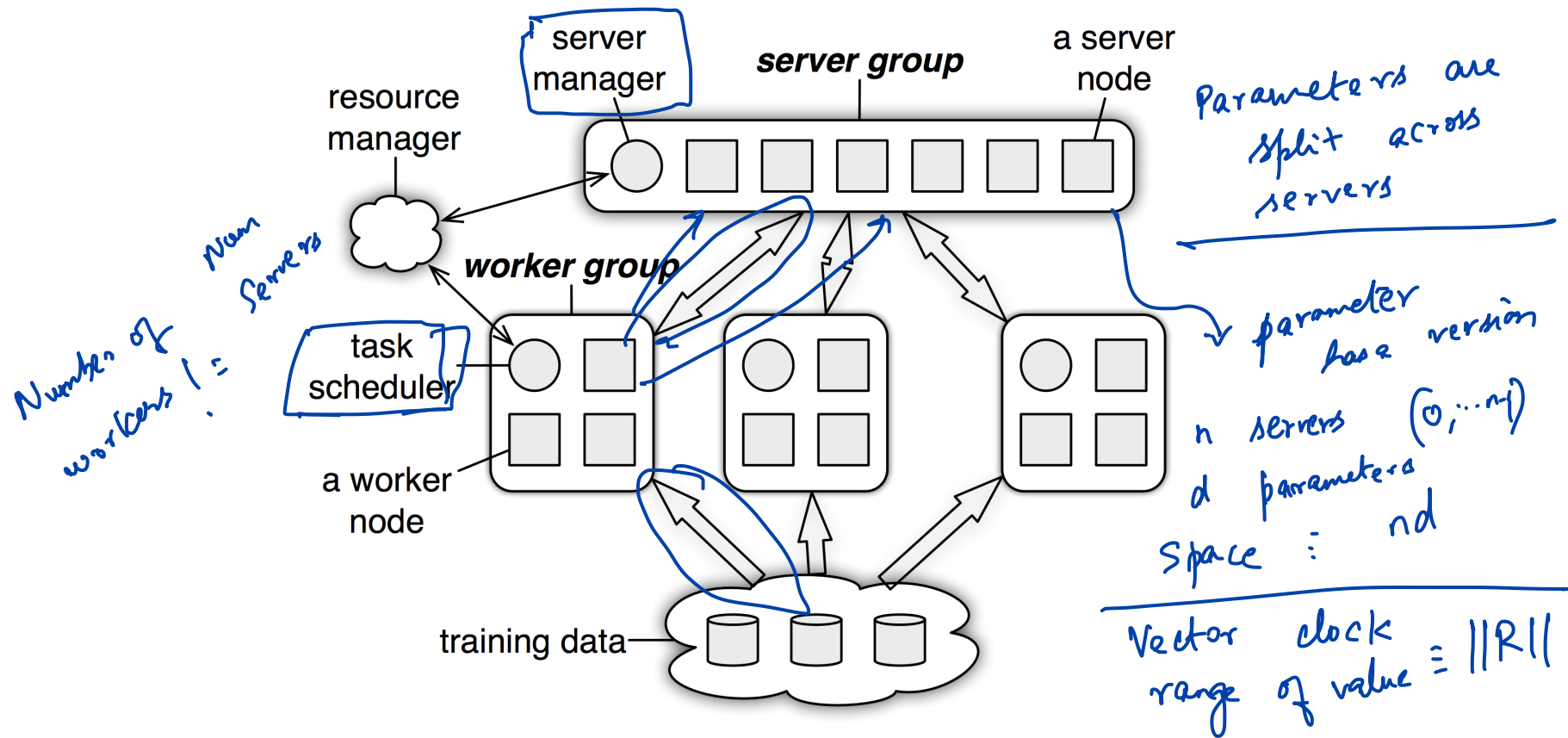
network traffic $\approx$ size of model

Computation flow

1. Sparse training data per worker
2. Compute $w^T x_i = g_i$
3. Update $w_2 = w_1 - \gamma \sum g_i$
4. Pull $w_2 \Rightarrow$ How long should we wait?



ARCHITECTURE



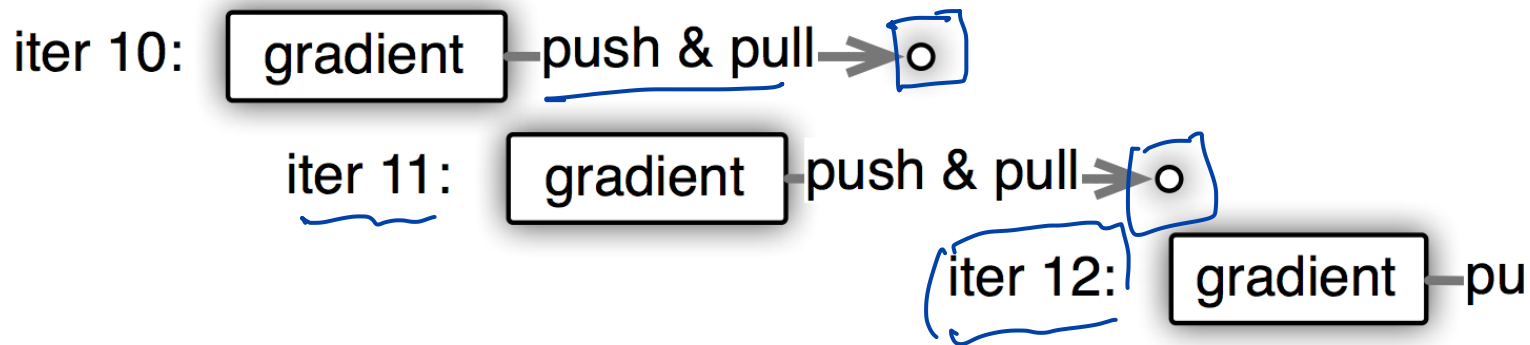
$$\begin{pmatrix} 0, 1, 5 \\ 1, 2, 5 \\ \vdots \end{pmatrix}$$

- $$\boxed{1.5 \mid 2.5 \mid 3.5 \mid \dots \dots \dots}$$

[illegible]

$w[R] = w[i=100]$ $\rightarrow$ destination

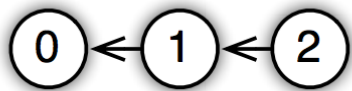
TASK DEPENDENCY



Workers have the option
of starting iteration with/without prior iteration

CONSISTENCY MODELS

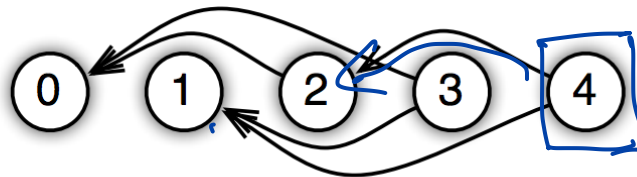
Flexible
Consistency



(a) Sequential



(b) Eventual



(c) 1 Bounded delay

BSP

User defined filters

Significantly modified filter

KKT filter

Not all
gradients need
to be transmitted

T_c Bounded

↳ Upper bound on
the staleness

$\nabla f : [\dots -1e-7 \dots 1e-10 \dots]$
million

IMPLEMENTATION: VECTOR CLOCKS

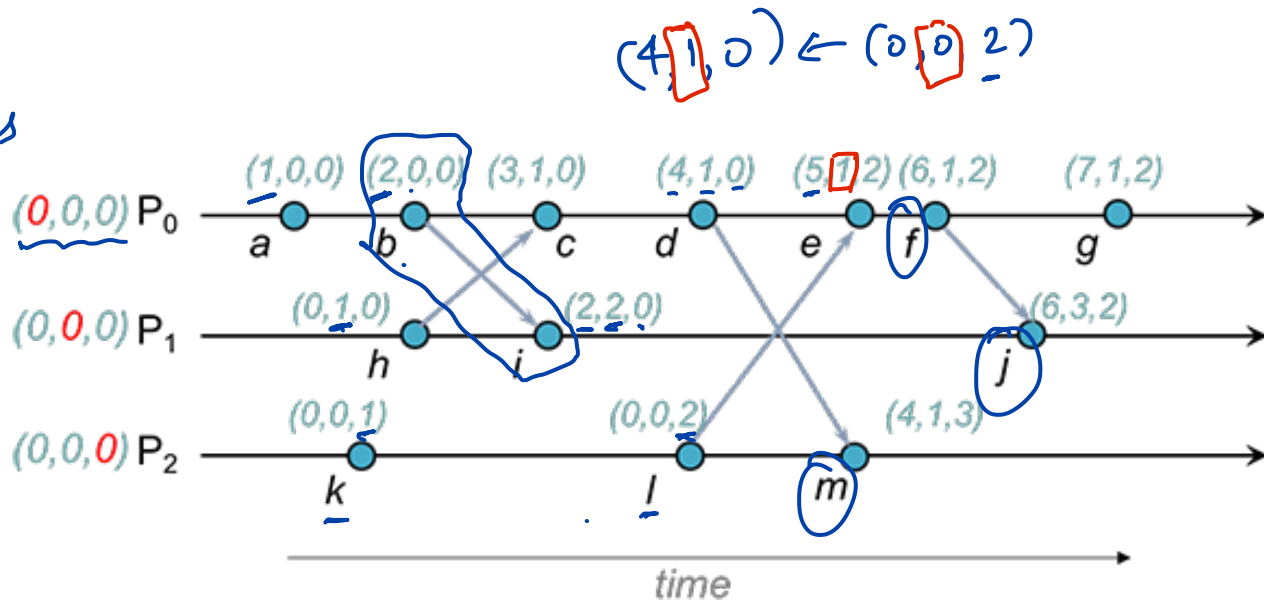
number of processes

Compute
→ inc local

send
→ inc local

recv
→ inc local

copy
counter
other
process
↳



if j before
or after
f?

$g = (7, 1, 2)$
 $j = (6, 3, 2)$ } Concurrent

$j: (6, 3, 2)$ f happens before j
 $f: (6, 1, 2)$

IMPLEMENTATION

Key Caching

- Worker might send the same key lists again
- Receiving node caches the key lists
- Sender only needs to send a hash of the list

0-100
255-302
⋮

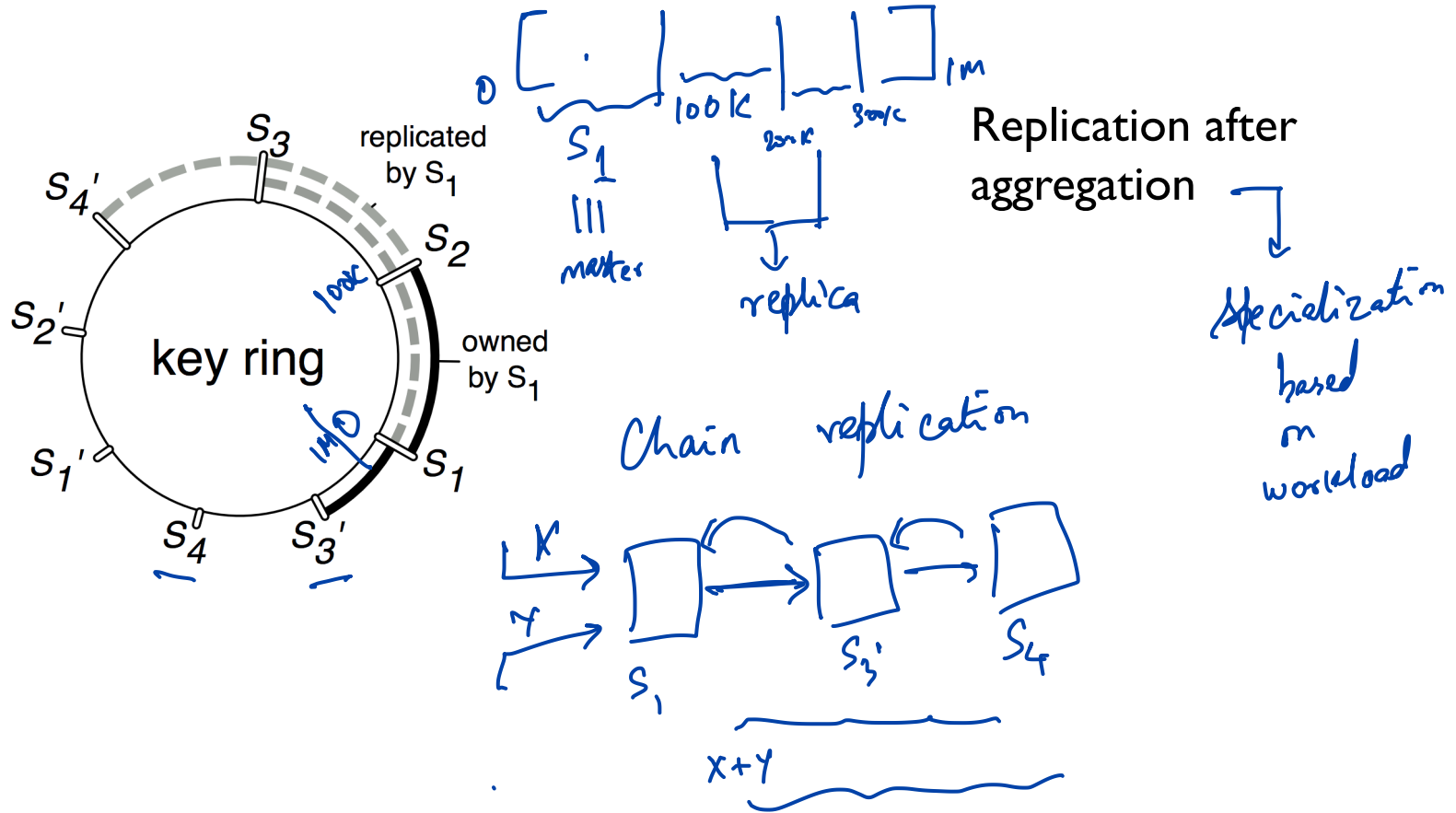
upto 2x

Value Compression

- Lots of repeated values, zeros
- Use Snappy to compress messages

1. Filter
2. values compress
3. Key Caching

IMPLEMENTATION: REPLICATION



FAULT TOLERANCE

1. Server manager assigns the new node a key range to serve as master.
2. The node fetches the range of data to maintain as master and k additional ranges to keep as slave.
3. The server manager broadcasts the node changes.
The recipients of the message may shrink their own data

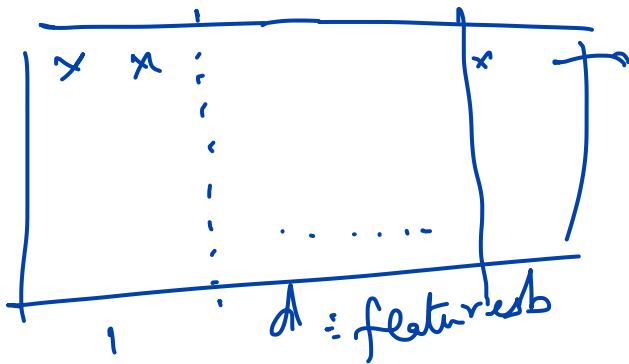
Worker node
fails

- Restart another worker
operates on same training data
- Ignore it
- Move training data to existing workers

Algorithm 3 Delayed Block Proximal Gradient [31]

Scheduler:

- 1: Partition features into b ranges $\mathcal{R}_1, \dots, \mathcal{R}_b$
- 2: **for** $t = 0$ **to** T **do**
- 3: Pick random range $\mathcal{R}_{i_t}$ and issue task to workers
- 4: **end for**



Worker r at iteration t

- 1: Wait until all iterations before $t - \tau$ are finished
- 2: Compute first-order gradient $g_r^{(t)}$ and diagonal second-order gradient $u_r^{(t)}$ on range $\mathcal{R}_{i_t}$
- 3: Push $g_r^{(t)}$ and $u_r^{(t)}$ to servers with the KKT filter
- 4: Pull $w_r^{(t+1)}$ from servers $\leftarrow$ block

Bounded Staleness

SPARSE LR

Servers at iteration t

- 1: Aggregate gradients to obtain $g^{(t)}$ and $u^{(t)}$
- 2: Solve the proximal operator

$$w^{(t+1)} \leftarrow \underset{u}{\operatorname{argmin}} \Omega(u) + \frac{1}{2\eta} \|w^{(t)} - \eta g^{(t)} + u\|_H^2,$$

where $H = \operatorname{diag}(h^{(t)})$ and $\|x\|_H^2 = x^T H x$

DISCUSSION

<https://forms.gle/35vrxyG6WLmSvCs38>

What are some of the downsides of using PS compared implementing Gradient Descent in Bismarck / Spark?

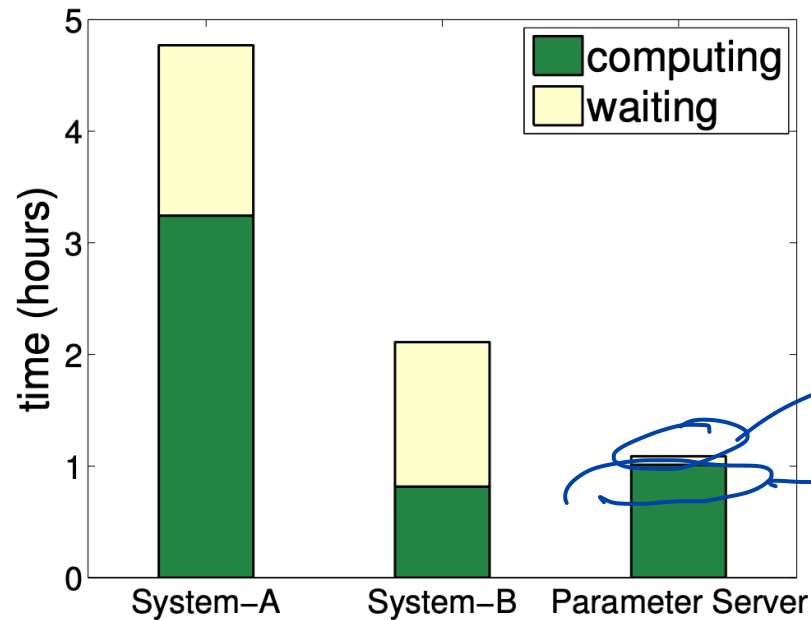
Training data fits in memory $\Rightarrow$ PS network overhead

Do we need full gradient? $\Rightarrow$ PS gradient is split.

Consistency (or lack of) ~~etc~~ lead to slow down

Sparse data PS is good $\Rightarrow$ What about dense data?

Extra resources to run servers



less waiting
more iterations
&
more compute

same algorithm

How would you integrate PS with a resource manager like Mesos? What would be some of the challenges?

↳ Allocate resources to run servers & workers

⇒ At least some fixed for servers

Not that flexible

↳ Long running workers & servers

min resource



pre-empt → time to converge

↳ Key cache change ⇒ training data changes

NEXT STEPS

Next class: Tensorflow

Assignment 2 is out!

Course project deadlines

- Oct 7 (topics)

- Oct 17 (proposals)