

Good morning!

CS 744: CLIPPER

Shivaram Venkataraman

Fall 2020

ADMINISTRIVIA

Course Project Proposals

- Due on Friday! →
- See Piazza for template
- Submission instructions soon

Midterm details → Oct 22 → Up to ML section

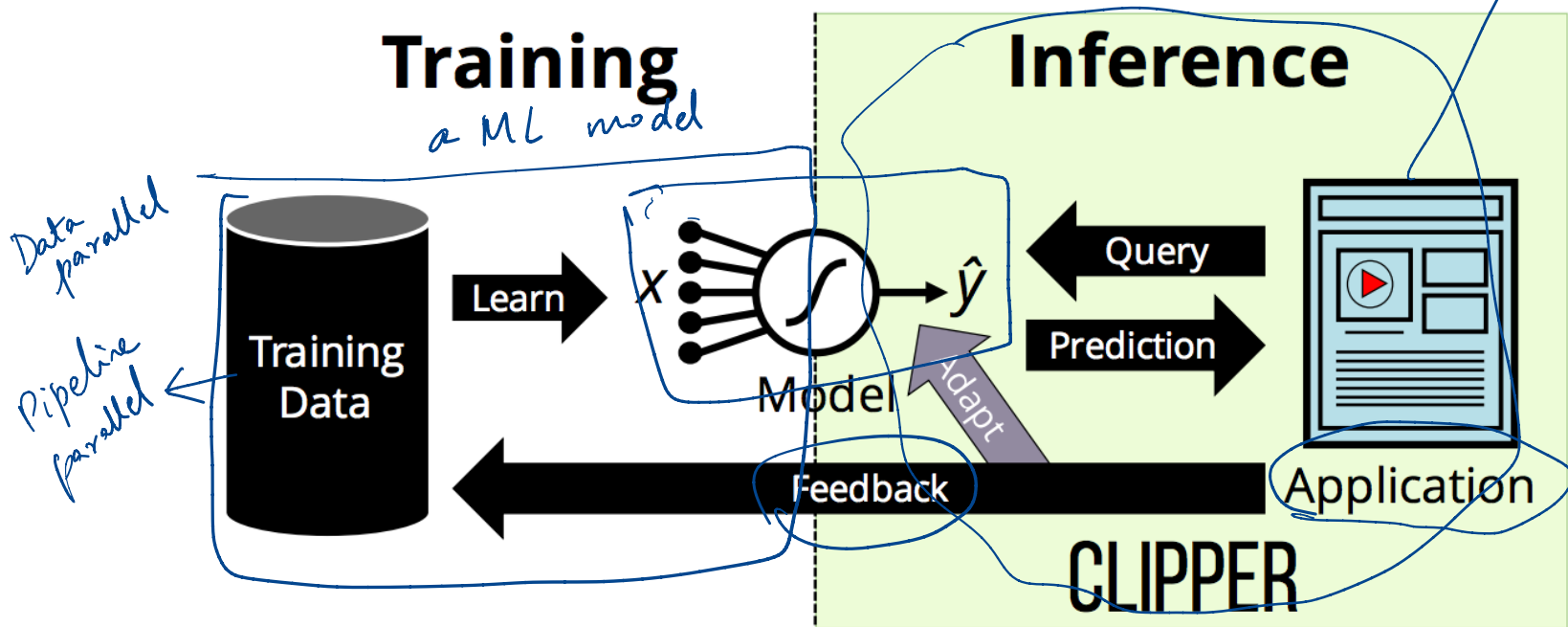
sample midterms
on Piazza

- Open book, open notes
- Held in class time 9.30-10.45am Central Time
- Type / Upload photos (extra 15 mins)

MACHINE LEARNING: INFERENCE

Reinforcement learning
↳ simulations

Netflix
recommend
movies



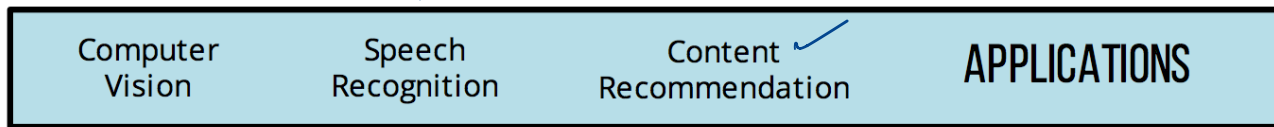
GOALS

- Interactive latencies (tail latency < 100ms) → 99 percentile or 99.9th percentile → low latency inference
- High throughput to handle load → many users or many requests that need to be made
- Improved prediction accuracy → ML specific
- Generality (?) → handle as many ML models / frameworks as possible

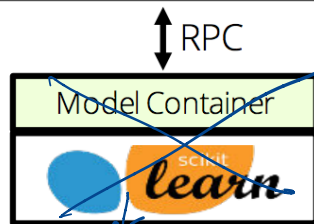
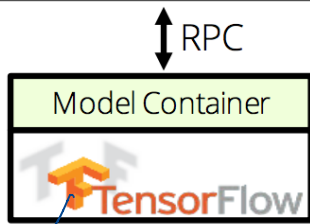
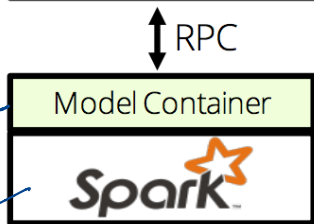
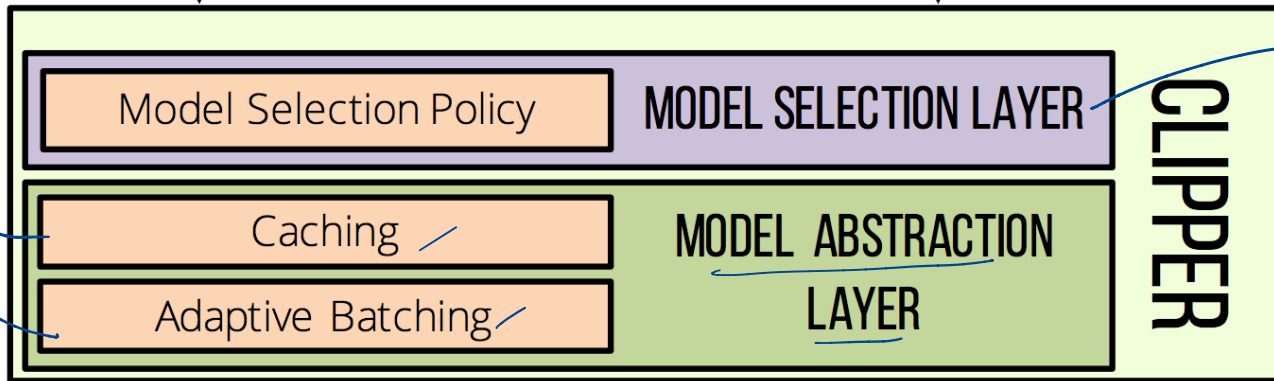
Tensorflow
Pytorch
MXnet
Scikit Learn

web server
↳ chat

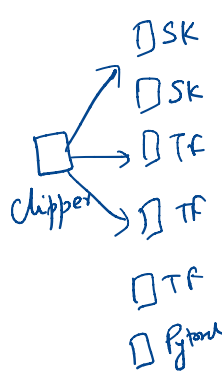
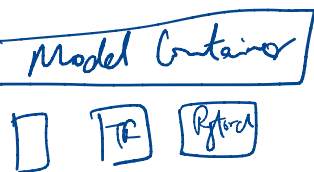
ARCHITECTURE



REST API



MACHINE LEARNING FRAMEWORKS



Requests over HTTP

Netflix

Improve accuracy

- Failures
- Replicates
 - ↳ 3 scikit
 - ↳ 1 spark

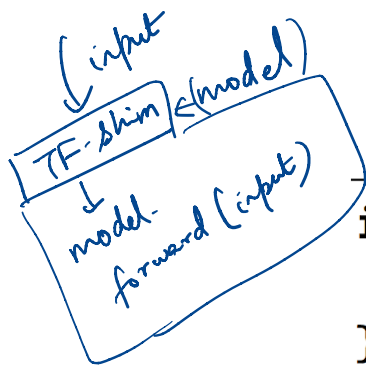
Higher throughput

Low latency

Generality goal

single machine

MODEL CONTAINERS



```
interface Predictor<X,Y> {  
    List<List<Y>> pred_batch(List<X> inputs);  
}
```

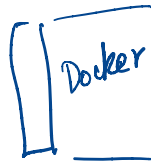
simple API

Interface is implemented
once per framework

Run using Docker containers

Can be replicated across machines

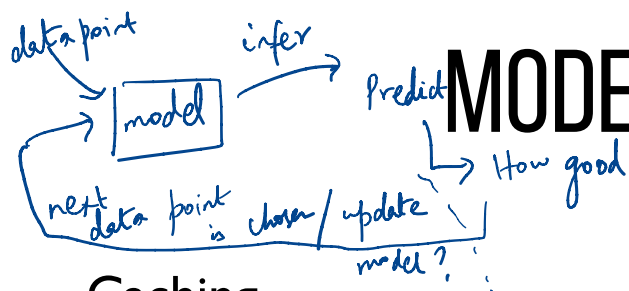
→ Model frameworks are
a black box
to clipper!



instantiate TF shim
with Resnet -18



instantiate TF shim
with Resnet -50



MODEL ABSTRACTION LAYER

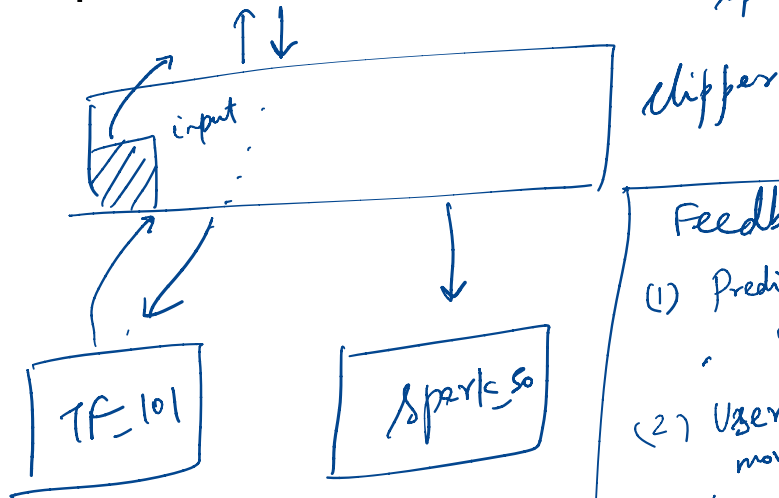
spec by provided the dev.

Caching

- Improve performance for frequent queries
- LRU eviction policy
- Important for feedback

high dim inputs
→ cache less hits likely! are

→ Predict (1) movies for user_id_1
TF_101 or Spark_50



Feedback

- (1) Predict Movies for user_1
- (2) User rates movie_1

Feedback and Predictions → Model update

To do an RPC
↳ fixed cost overhead
→ amortize

BATCHING, QUEUING

batch size that max.
tput while within
latency SLO,

each model container
has diff overheads,
parallelism.
Optimal batch size
for each model! could vary

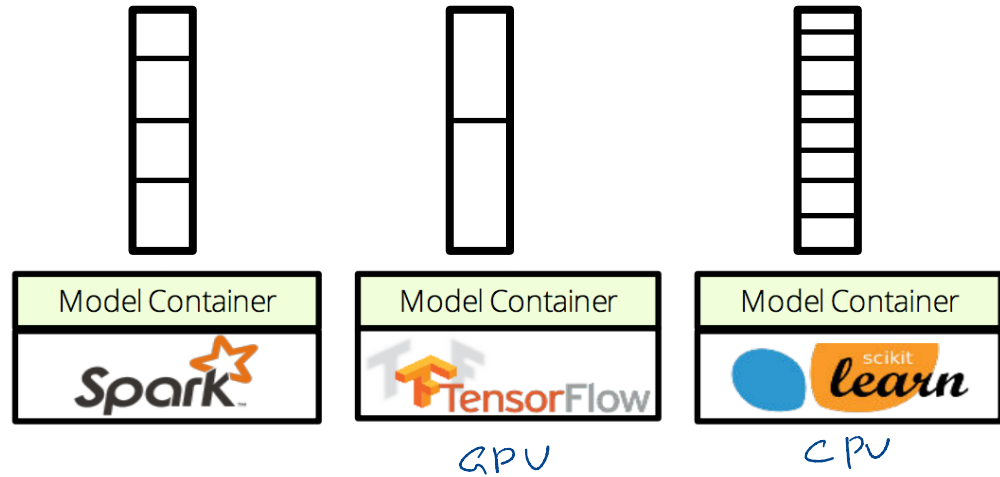
batches use
hardware
parallelism

Goals, Insight

- Increase latency (within SLO) for improved throughput
- Reduce RPC overheads
- GPU / BLAS acceleration

Approach

- Per container queues.
- Why?



ADAPTIVE BATCHING

AIMD: Additive Inc Multiplicative Dec (AIMD)

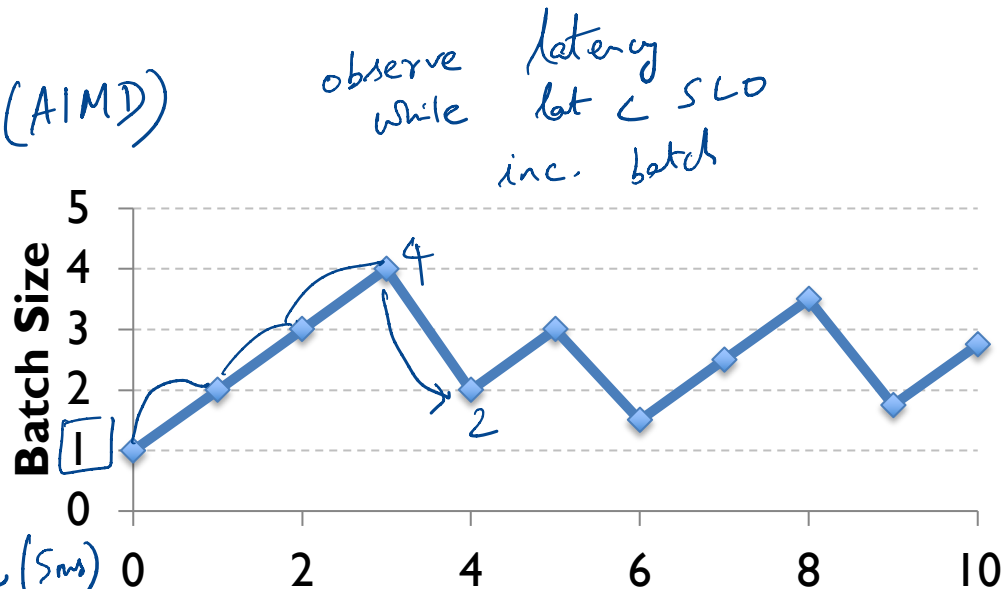
Why?

Increase batch size carefully
& Latency quickly comes down if it's above SLO!

Delayed: Wait until batch exists

Why?

Collect examples
up to a certain
time then dispatch



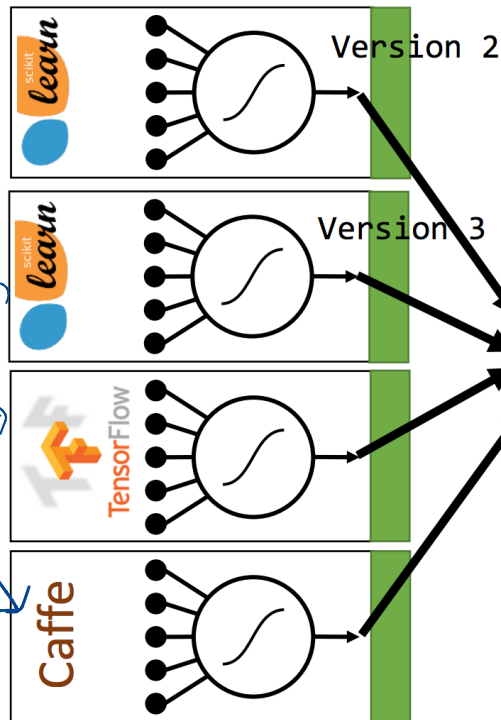
Time (inc. lat) should I wait?
(dec. util) should I send?



MODEL SELECTION

→ Improve Accuracy

ensembles



“CAT”
“CAT”
“CAT”
“CAT”

Policy

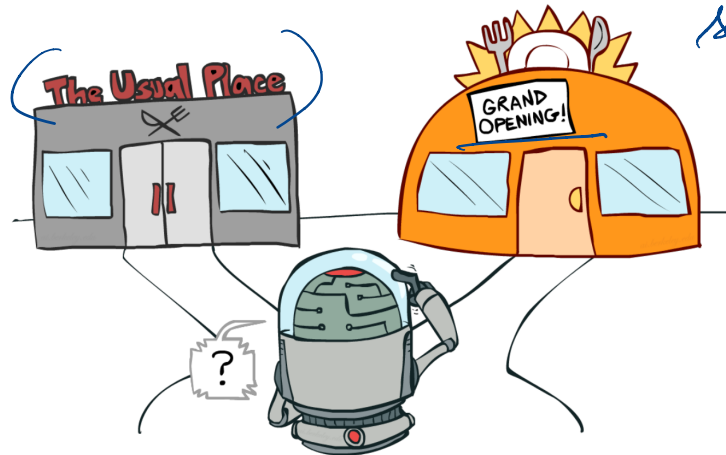
“CAT”
CONFIDENT

SINGLE MODEL SELECTION

Given you have N models
↳ which one should I use?

Multi-Arm Bandit formulation

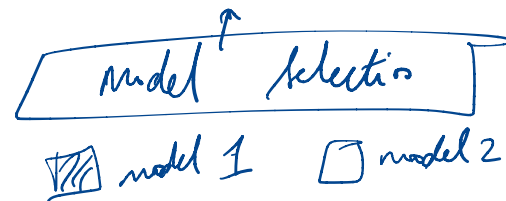
- Explore vs Exploit
- Regret: Loss by not picking optimal action
- Goal: Minimize regret



Clipper

- Exp3 algorithm
- Single evaluation
- Scales to more models
- update weights based on feedback

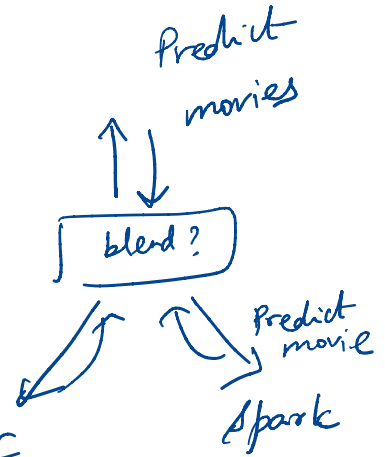
associate a weight with each option



MULTI MODELS → ensembles

Ensemble

- Combine output from models (weighted average)
- How do we get the weights ?

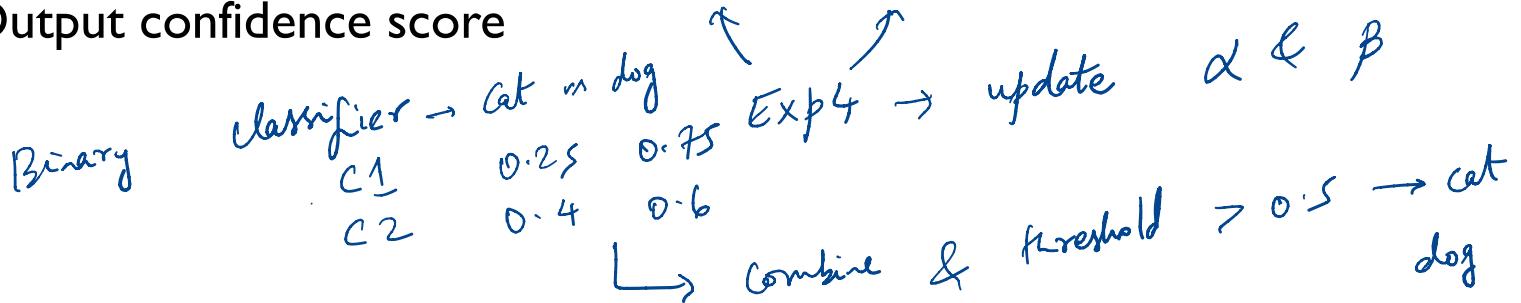


Robust Prediction

- React to model changes
- Output confidence score

linear combination TF

$$= \alpha y_i + \beta z_i$$



STRAGGLER MITIGATION

Why do stragglers occur?

99th p_c latency

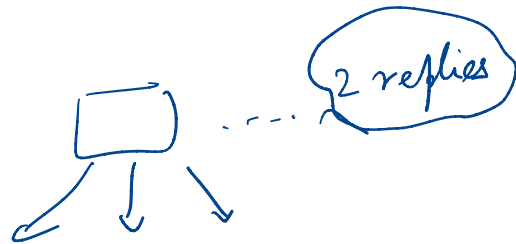
↳ If we wait for N model containers to reply, some of them might be slow?

↳ more replicas / scaling?

Approach

↳ Approx result based on whatever has finished

→ Better approx than late! → ML specific



SUMMARY

- Clipper: ML inference Workloads + Requirements
- Layered architecture provides **generality**
- Caching, Batching, Replication to improve **latency, throughput**
- Multi-Arm bandits to improve **accuracy**

DISCUSSION

<https://forms.gle/FCVhPURqz7HSbDtg6>

Consider a scenario where you run a model serving service that hosts a number of different applications. The traffic for some applications is sporadic (e.g. only a few hours where they are used). What are some advantages / disadvantages of using Clipper for such a service?

Advantages

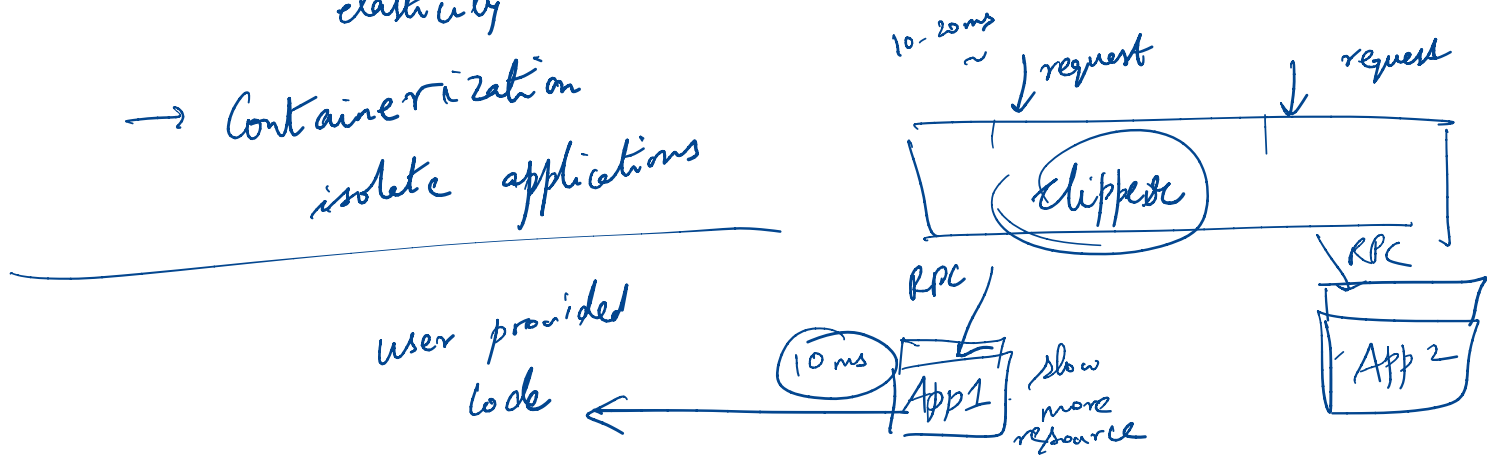
→ Adaptive batching
+ delayed → tune

→ multiple replicas
elasticity

→ Containerization
isolate applications

Disadvantages
→ Cache might be contended

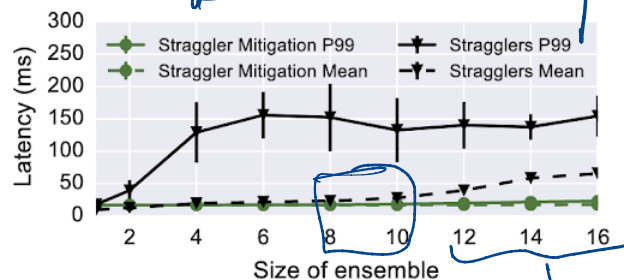
→ how to do this
online & fast?



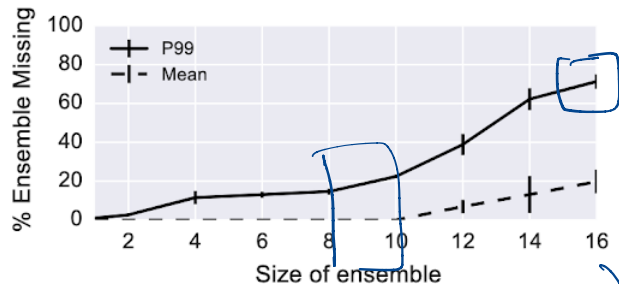
Larger ensembles use a lot of resources

8-10 seems to a good size - balancing bunch of different things?

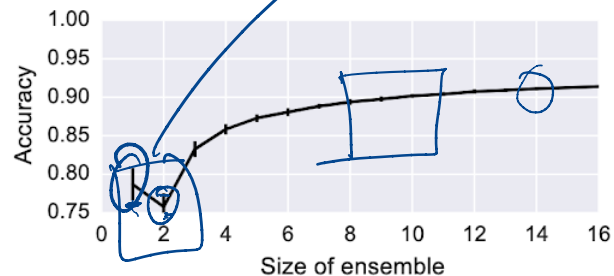
users going on here?



(a) Latency



(b) Missing Predictions



(c) Accuracy

16 ensembles latency inflation is very low

reasonable accuracy