

Learning to Solve Routing Problems



Wouter Kool



Herke van Hoof



Joaquim Gromicho



Max Welling



Wouter

Master of Business Analytics &
Master of Econometrics & OR



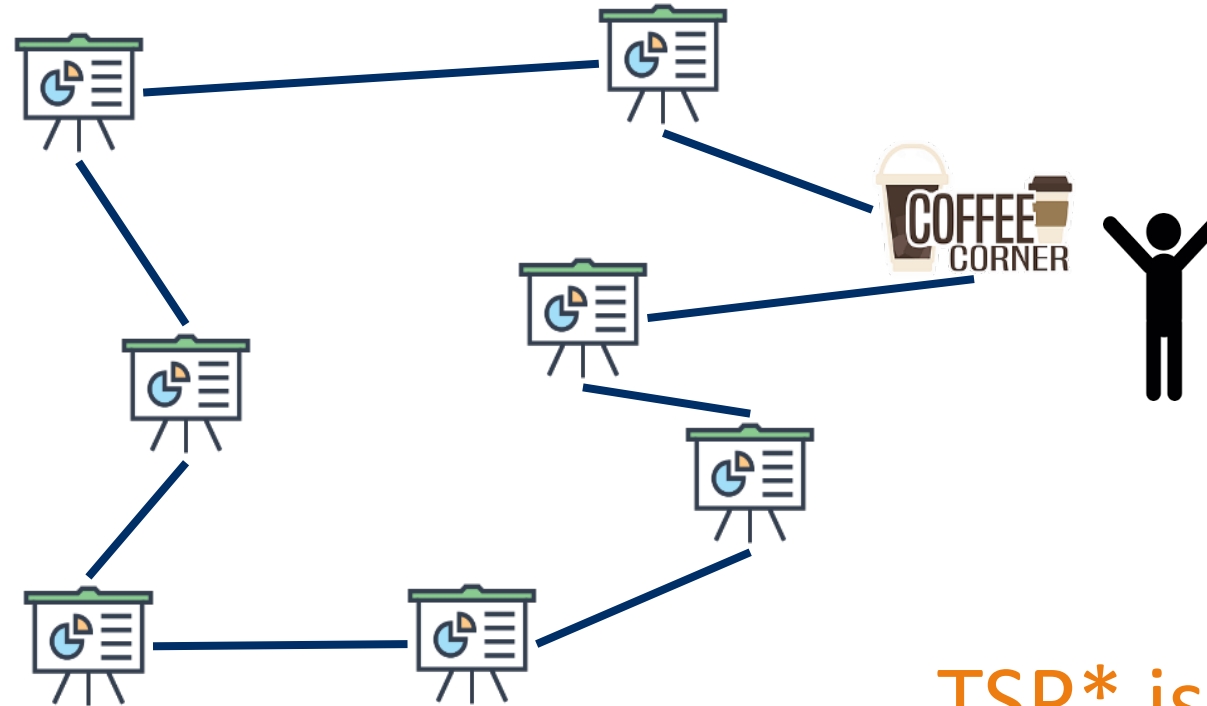
1+6 years at **ORTEC**, OR engineer

September 2017, PhD @

Working on ML & OR



Travelling Scientist Problem (TSP)



TSP* is (NP-)hard!

Kool et al., 2019

What does it mean?

Finding *optimal* solutions for *all* problem instances

Finding *acceptable* solutions for *relevant* problem instances

MISSION:
IMPOSSIBLE

* unless $P = NP$

MISSION:
~~IMPOSSIBLE~~

We use HEURISTICS
Can be seen as 'rules of thumb'

'next location should be nearby'

Designing of heuristics is like feature engineering

**HARD
WORK**

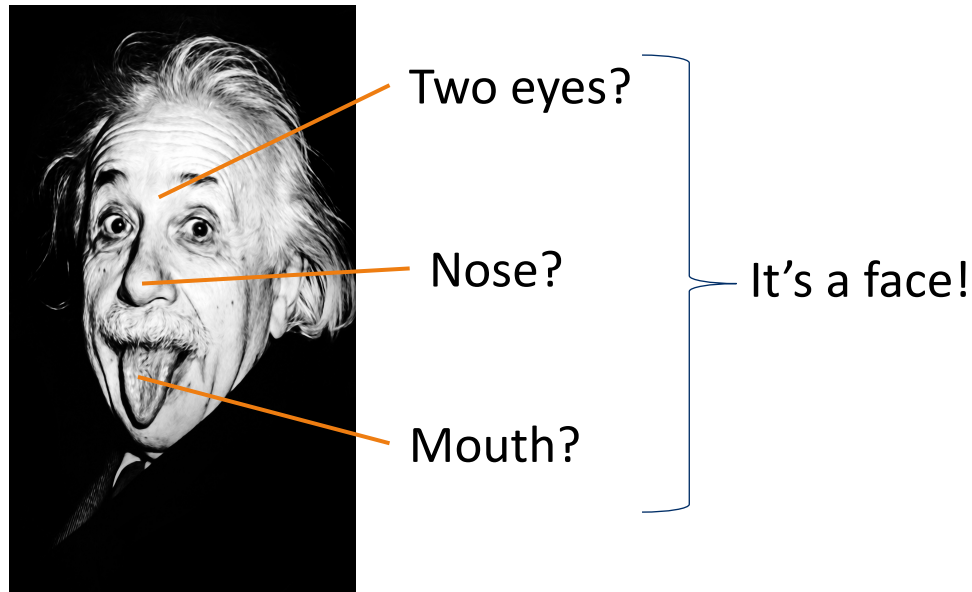
Computer Vision Features
(SIFT, etc.)

Feature engineering

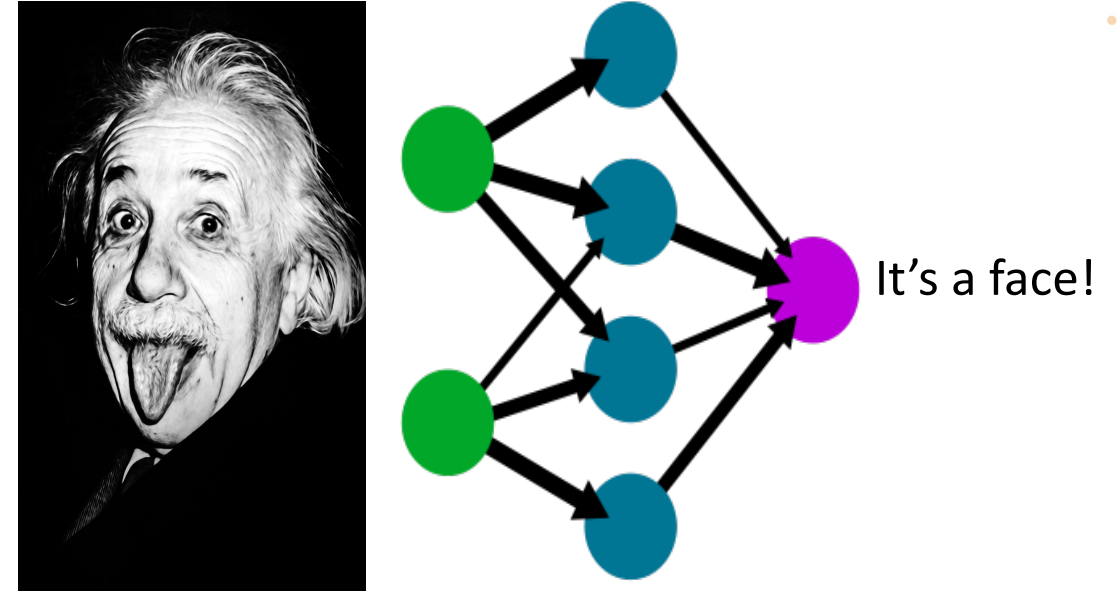
- Needs expert knowledge
- Time consuming hand-tuning

So what do we do?

Designing of heuristics is like feature engineering

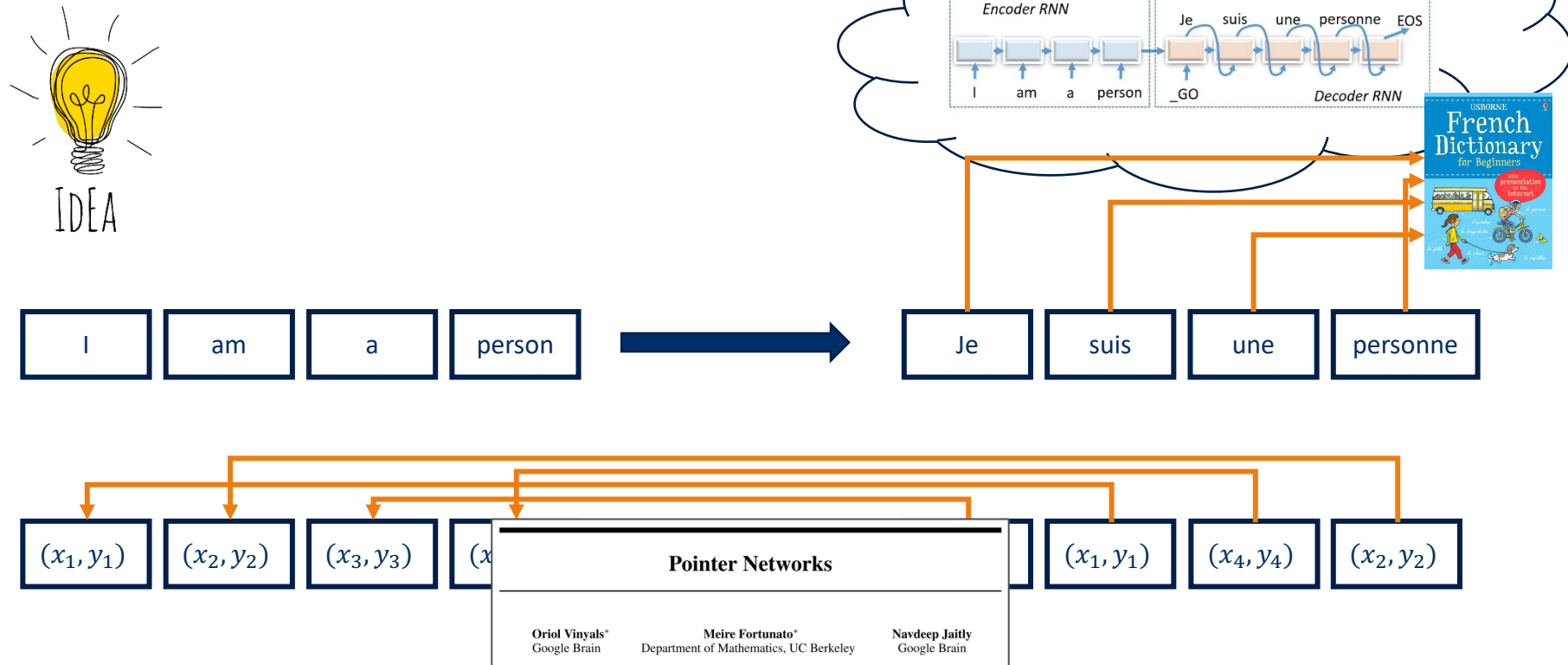


Traditional approach
Feature engineering



Deep Learning
No feature engineering

'Translate' problem into solution

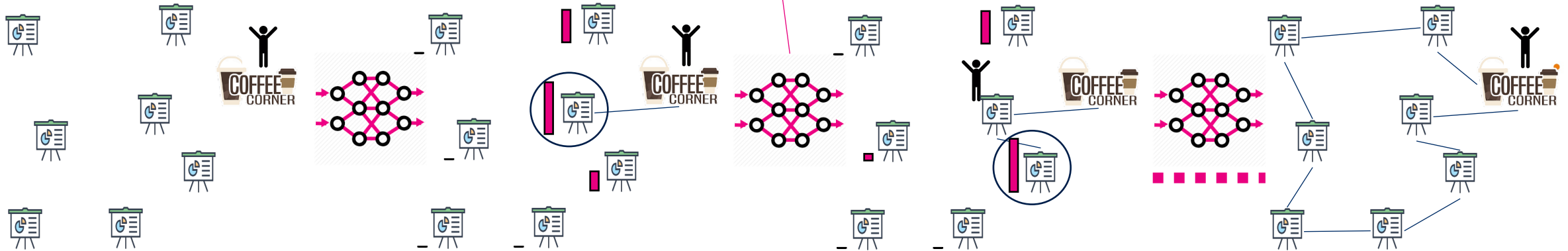


How does that work?

Instance $s = ((x_1, y_1), (x_2, y_2), \dots, (x_n, y_n))$

Model $p_{\theta}(\pi_t | s, \pi_{<t}) = p_{\theta}(\text{next node} \mid \text{partial tour})$

Solution $\pi = (\pi_1, \pi_2, \dots)$ with length $L(\pi)$



Sample $\pi_1 \sim p_{\theta}(\pi_1 | s)$

Sample $\pi_2 \sim p_{\theta}(\pi_2 | s, \pi_1)$

Sample $\pi_t \sim p_{\theta}(\pi_t | s, \pi_{<t})$

Randomized algorithm
with expected cost:

$$E_{p_{\theta}(\pi | s)} [L(\pi)]$$

How to optimize θ ?

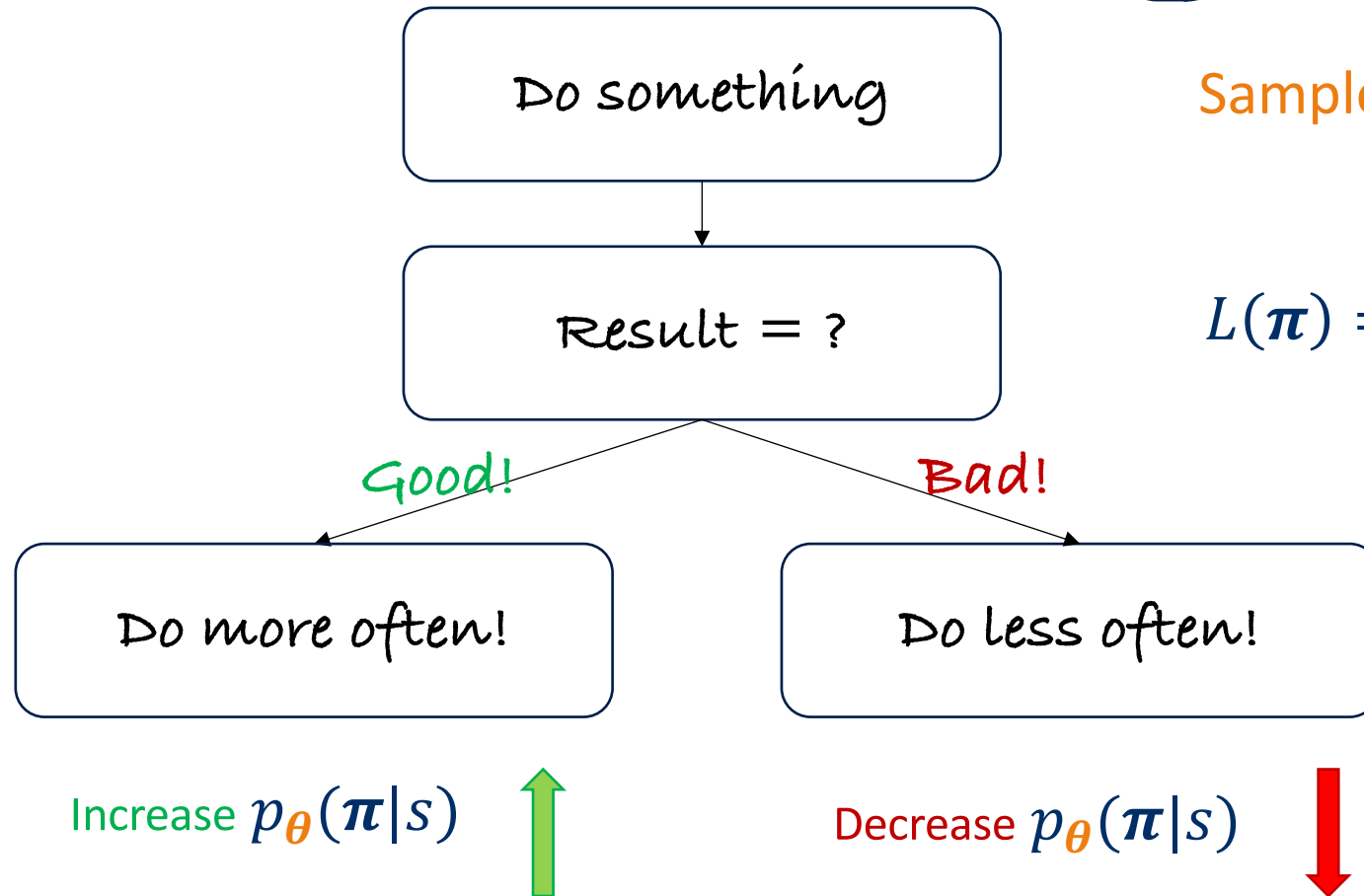
NEURAL COMBINATORIAL OPTIMIZATION
WITH REINFORCEMENT LEARNING

Irwan Bello*, Hieu Pham*, Quoc V. Le, Mohammad Norouzi, Samy Bengio
Google Brain
{ibello, hyhieu, qvl, mnorouzi, bengio}@google.com

REINFORCE (for dummies)



Repeat



Sample $\pi \sim p_{\theta}(\cdot | s)$

$L(\pi) = 7.43$

We need a *baseline* to compare against:
rollout earlier model

$\pi^{bl} \sim p_{\theta^{bl}}(\cdot | s)$ (greedy!)

$L(\pi^{bl}) = 6.89$

The rollout baseline



Use (rollout) the model but *greedy* instead of sampling!

Sample $\pi \sim p_{\theta}(\cdot | s)$

Rollout $\pi^{bl} \sim p_{\theta^{bl}}(\cdot | s)$ (greedy!)

$L(\pi) < L(\pi^{bl})$ Good!

$L(\pi) > L(\pi^{bl})$ Bad!

Adjust $p_{\theta}(\pi | s)$
proportional to
 $L(\pi) - L(\pi^{bl})$

Learning Algorithm (roughly)

Init $\theta, \theta^{bl} \leftarrow \theta$

For ever(y epoch):

For iteration:

Sample s

Sample $\pi \sim p_{\theta}(\cdot | s)$

Rollout $\pi^{bl} \sim p_{\theta^{bl}}(\cdot | s)$ (greedy!)

Update $\theta \leftarrow \theta - \eta \nabla \log p_{\theta}(\pi | s) (L(\pi) - L(\pi^{bl}))$

If θ better* than θ^{bl}

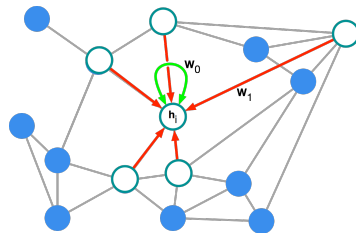
Update $\theta^{bl} \leftarrow \theta$

* Paired t -test on solution of 10 000 instances with greedy rollout

What's the model architecture?

$$p_{\theta}(\pi_t | s, \pi_{<t})$$

Graph convolutions



+

Attention Is All You Need

Ashish Vaswani*
Google Brain
avaswani@google.com

Noam Shazeer*
Google Brain
noam@google.com

Niki Parmar*
Google Research
nikip@google.com

Jakob Uszkoreit*
Google Research
usz@google.com

Llion Jones*
Google Research
llion@google.com

Aidan N. Gomez*[†]
University of Toronto
aidan@cs.toronto.edu

Lukasz Kaiser*
Google Brain
lukasz@kaiser@google.com

Illia Polosukhin*[†]
illia.polosukhin@gmail.com

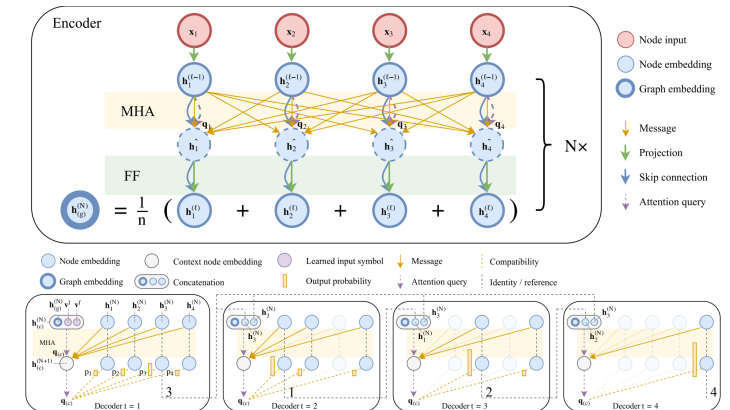
Read the paper...

ATTENTION, LEARN TO SOLVE ROUTING PROBLEMS!

Wouter Kool
University of Amsterdam
w.w.m.kool@uva.nl

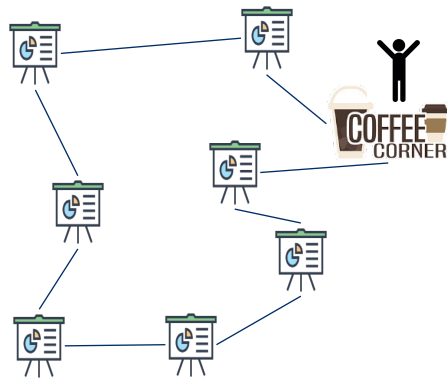
Herke van Hoof
University of Amsterdam
h.c.vanhoof@uva.nl

Max Welling
University of Amsterdam
m.welling@uva.nl



Experiments

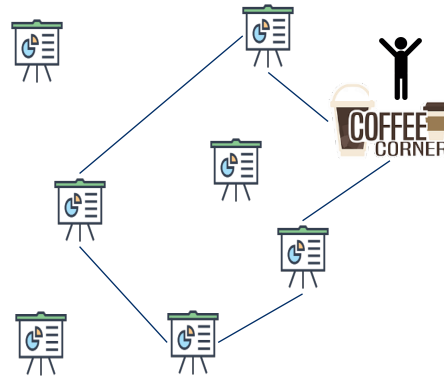
Travelling Salesman Problem (TSP)



Minimize length

Visit all nodes

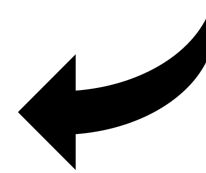
Orienteering Problem (OP)



Maximize total prize

Max length constraint

(Stochastic) Prize Collecting TSP ((s)PCTSP)

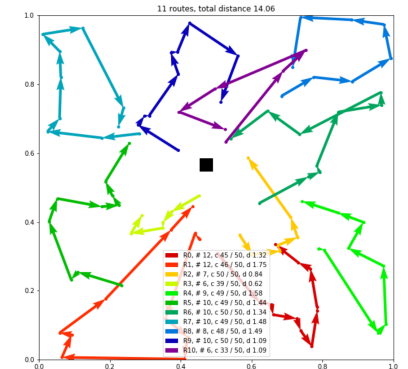


Minimize length + penalties

of unvisited nodes

Collect minimum total prize

Vehicle Routing Problem (VRP)



Minimize length

Visit all nodes

Total route demand must fit vehicle capacity

Train for each problem, *same hyperparameters!*

Results Attention Model + Rollout Baseline

- Improves over classical heuristics!
- Improves over prior learned heuristics!
 - Attention Model improves
 - Rollout helps significantly
- Gets close to single-purpose SOTA (20 to 100 nodes!)
 - TSP 0.34% to 4.53% (greedy)
 - TSP 0.08% to 2.26% (best of 1280 samples)

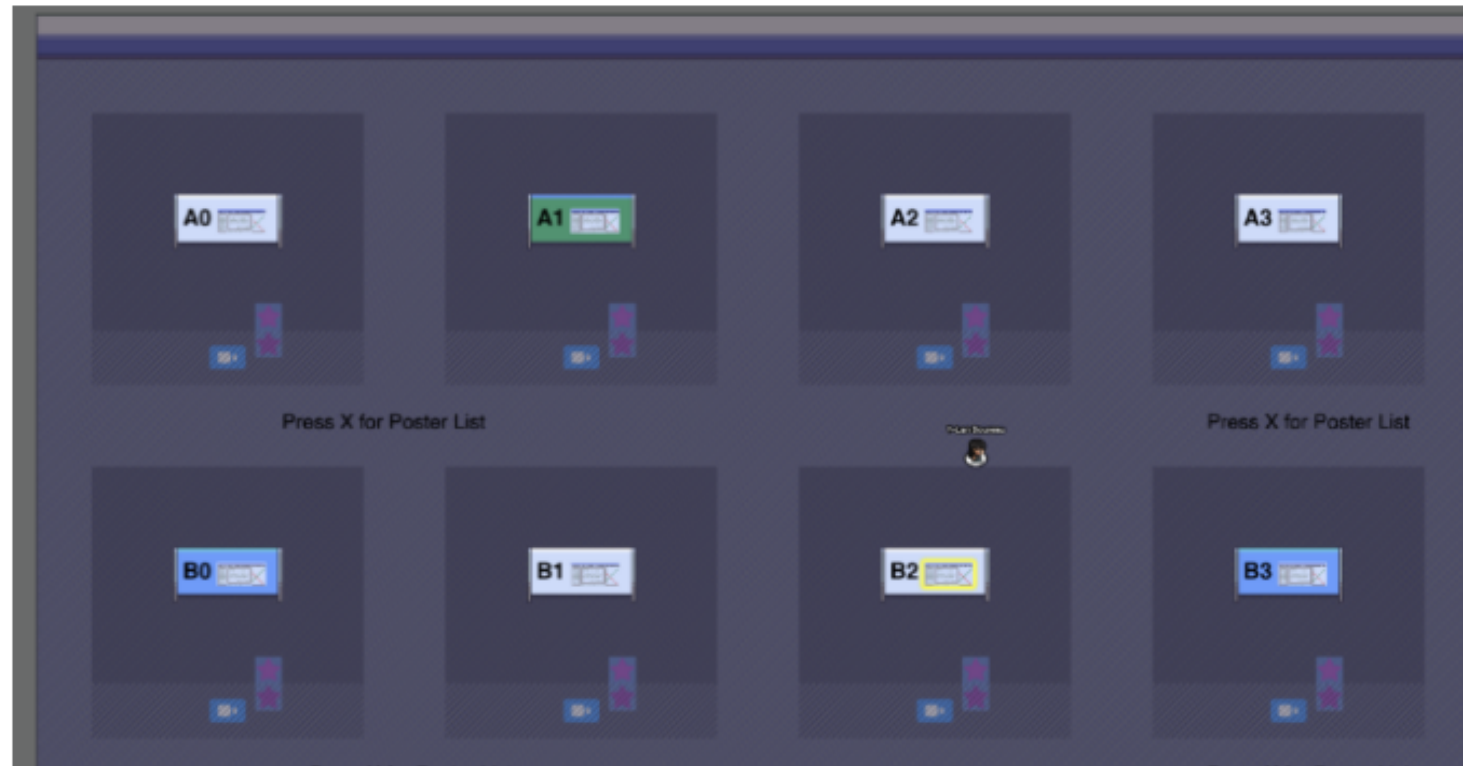
Neural Information Processing Systems Conference

Tweets sent to this account are not actively monitored. To contact us please go to <http://neurips.cc/Help/Contact>

Follow



simply wandering; each poster had a clearly marked presenter spot to easily spot the presenter; people could teleport directly to the poster of their choice from the NeurIPS website, and a coordinate systems allowed people to locate a poster of interest once they were in a room.

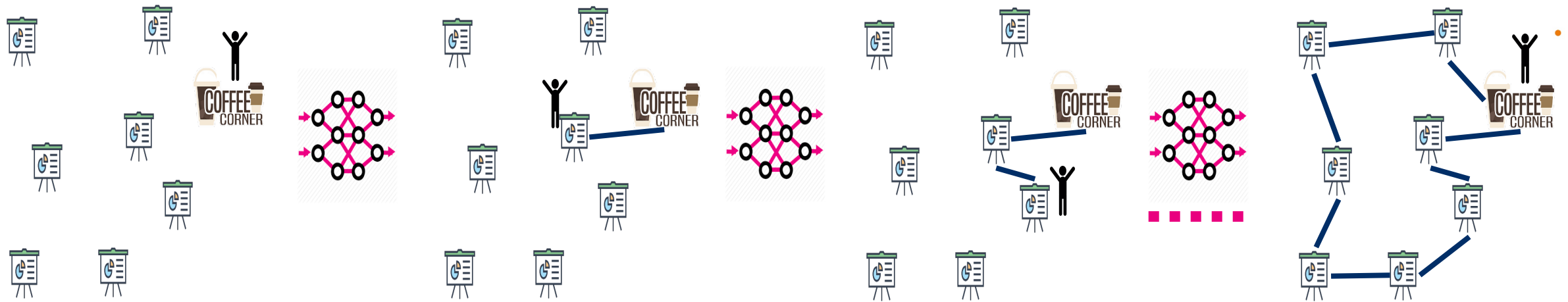


<https://neuripsconf.medium.com/neurips-2020-online-experiments-gather-town-poster-sessions-and-mementor-ac1573d61c8a>

Something else!

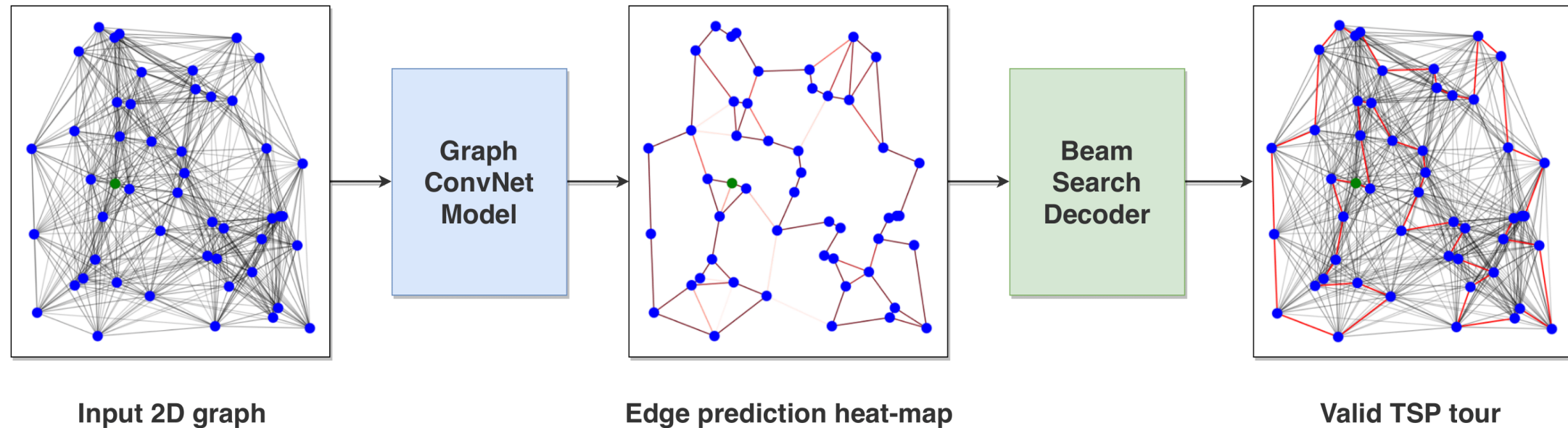
Deep Policy Dynamic Programming

Autoregressive approach



Vinyals et al., 2015
Bello et al., 2016
Kool et al., 2019

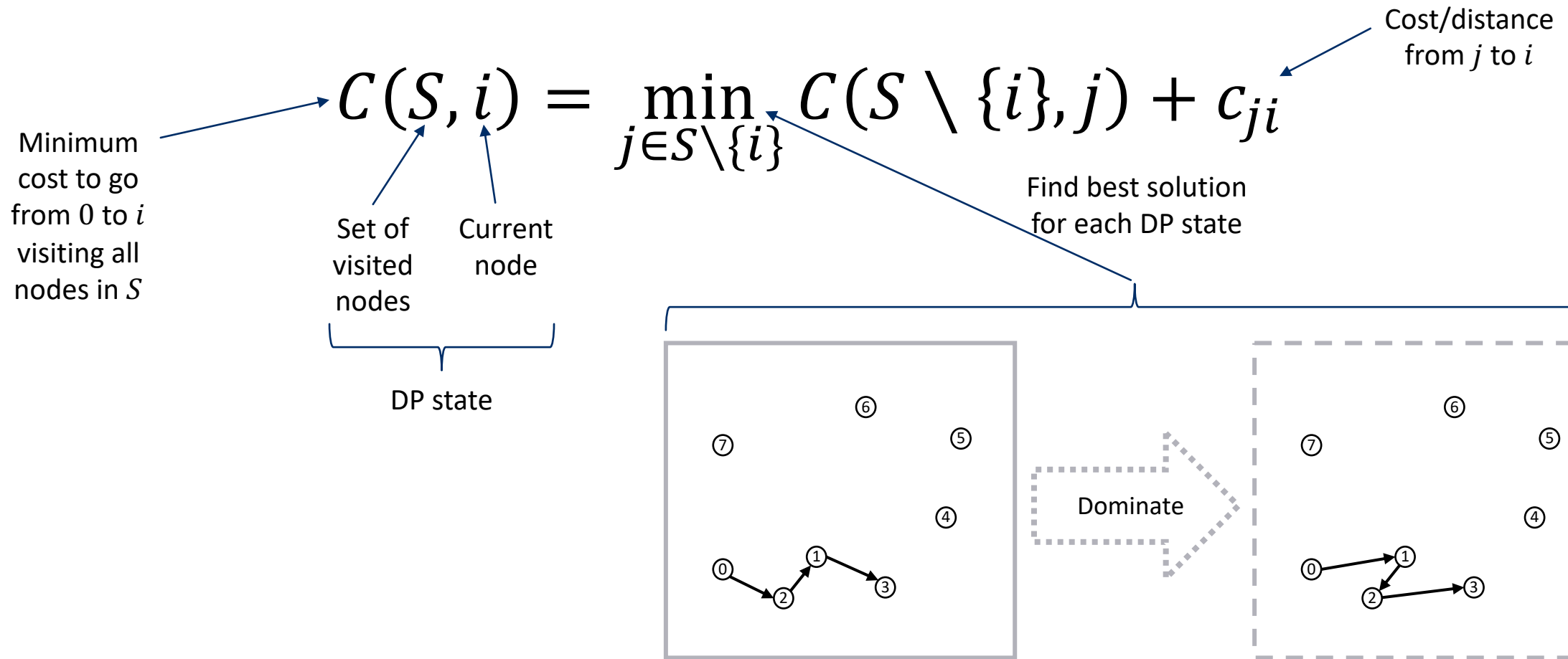
Non-autoregressive approach (Joshi et al., 2019)



Picture by Joshi et al., 2019

Note: trained using supervised learning!

Dynamic Programming for TSP (Held & Karp, 1962; Bellman, 1962)

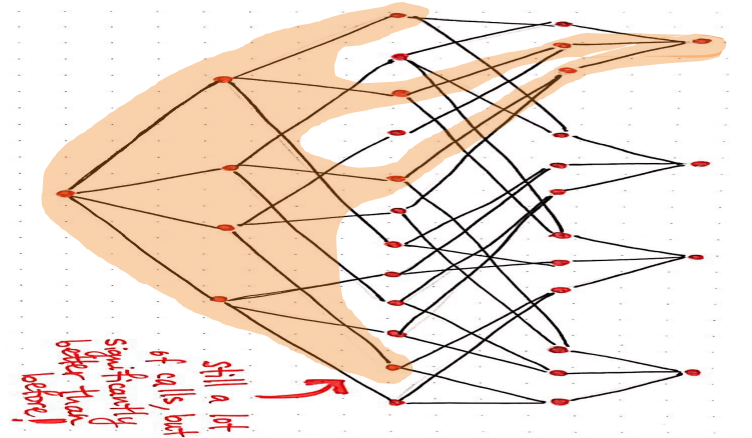


Deep Policy Dynamic Programming (DPDP)

<https://arxiv.org/abs/2102.11756>

DPDP

$O(Bn)$ or linear



For each iteration

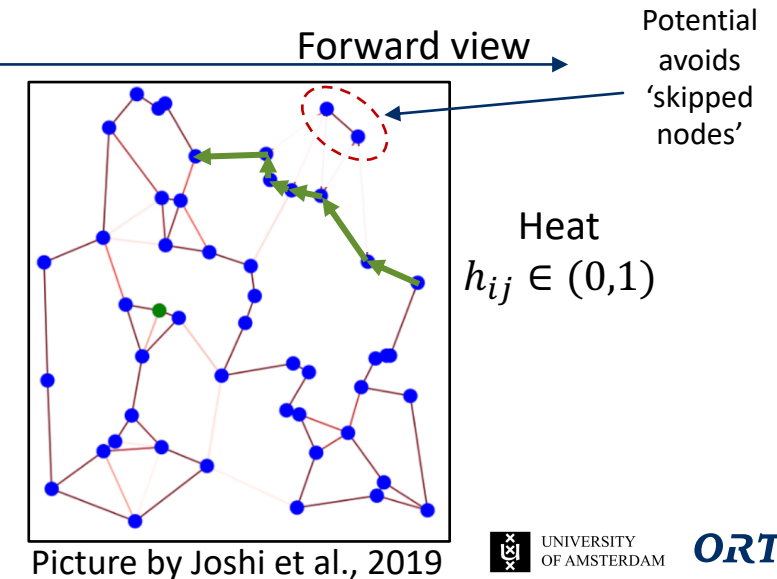
- Expand solutions
- Remove dominated solutions
- Select top B according to *policy*
- Repeat

Policy: select top B solutions that maximize the score.

$$\text{SCORE} = \text{HEAT} + \text{POTENTIAL}$$

Heat of edges
in solution

Estimate for
unvisited nodes
based on remaining edges



Deep Policy Dynamic Programming (DPDP)

<https://arxiv.org/abs/2102.11756>

- DPDP is a *beam search* over the DP state space, guided by a neural network
- DP is flexible framework for many VRP variants e.g. time windows
- Suitable for GPU implementation
- Natural trade-off compute vs. performance -> asymptotically optimal
- Supervised training based on example solutions
- Test time: only evaluate NN once!

Results

Travelling Salesman Problem

Table 1. Main results for TSP with 100 nodes.

METHOD	COST	GAP	TIME
CONCORDE	7.765	0.000 %	6M
LKH	7.765	0.000 %	42M
GUROBI	7.776	0.15 %	31M
KOOL ET AL. (2019)	7.94	2.26 %	1H
JOSHI ET AL. (2019A)	7.87	1.39 %	40M
DA COSTA ET AL. (2020)	7.83	0.87 %	41M
FU ET AL. (2020)	7.764*	0.04 %	4M + 11M
DPDP 10K	7.765	0.009 %	10M + 1H06M
DPDP 100K	7.765	0.004 %	10M + 2H35M

Vehicle Routing Problem

Table 2. Main results for VRP with 100 nodes.

METHOD	COST	GAP	TIME
LKH	15.647	0.000 %	12H59M
XIN ET AL. (2020)	16.49	4.99 %	39s
KOOL ET AL. (2019)	16.23	3.72 %	2H
CHEN & TIAN (2019)	16.10	2.90 %	1H
PENG ET AL. (2019)	16.27	3.96 %	6H
WU ET AL. (2019)	16.03	2.47 %	5H
HOTTUNG & TIERNEY (2019)	15.99	1.02 %	1H
LU ET AL. (2020)	15.57*	-	4000H
DPDP 10K	15.832	1.183 %	10M + 2H24M
DPDP 100K	15.694	0.305 %	10M + 5H48M
DPDP 1M	15.627	- 0.127 %	10M + 48H27M

The end

Questions?