



GNNLab: A Factored System for *Sample-based GNN Training* over GPUs

RONG CHEN

IPADS, SJTU

ASAIL 2022

Joint work with Jianbang, Dahai, Xiaoni @IPADS, Qiang @BASICS, and Lei, Wenyan @Ailibaba

Who AM I



2

Rong Chen (陈榕) / SJTU



https://ipads.se.sjtu.edu.cn/rong_chen



- ▶ Institute of Parallel And Distributed Systems (IPADS)
- ▶ **Best paper award:** EuroSys 2015, APsys 2017, and ICPP 2007
- ▶ 7 OSDI/SOSP papers and 9 EuroSys/Usenix ATC papers
- ▶ Huawei OlympusMons Pioneer Award, 2020

Research Interest

- ▶ Building *efficient*, scalable and reliable distributed framework

GNN Training



Whole-graph training: hard to scale

Whole-graph training: hard to scale

Sample-based training

- ▶ Systems: DGL, PyG, AliGraph^[VLDB'19], P3^[OSDI'21], . . . , GNNLab
- ▶ **Friendly to GPU:** massive parallelisms and limited GPU memory
- ▶ SET model: Sample, Extract and Train

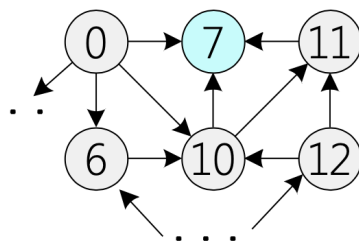
Sample-base GNN Training



5

SET model

Graph topc



Features

0	..	6	7	..	10	11	12
A ₁	..	G ₁	H ₁	..	K ₁	L ₁	M ₁
A ₂	..	G ₂	H ₂	..	K ₂	L ₂	M ₂
A ₃	..	G ₃	H ₃	..	K ₃	L ₃	M ₃
A ₄	..	G ₄	H ₄	..	K ₄	L ₄	M ₄

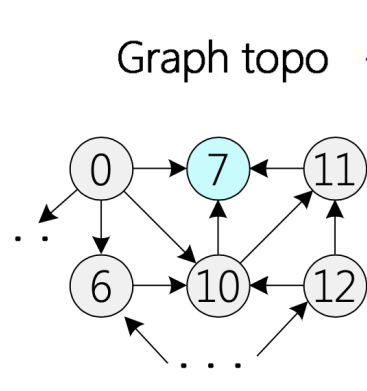
Sample-base GNN Training



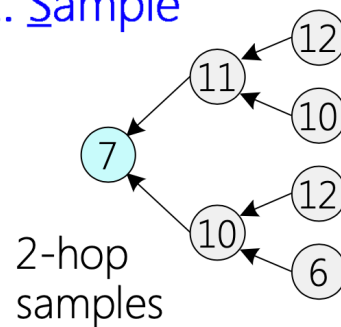
6

SET model

1. Sample

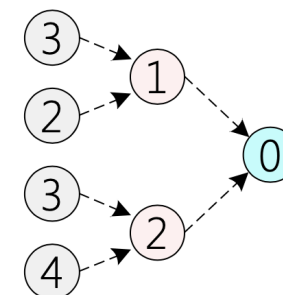


1. Sample



Vertex ID
mapping

7	0
11	1
10	2
12	3
6	4



Features

	0	..	6	7	..	10	11	12
A ₁	..		G ₁	H ₁	..	K ₁	L ₁	M ₁
A ₂	..		G ₂	H ₂	..	K ₂	L ₂	M ₂
A ₃	..		G ₃	H ₃	..	K ₃	L ₃	M ₃
A ₄	..		G ₄	H ₄	..	K ₄	L ₄	M ₄

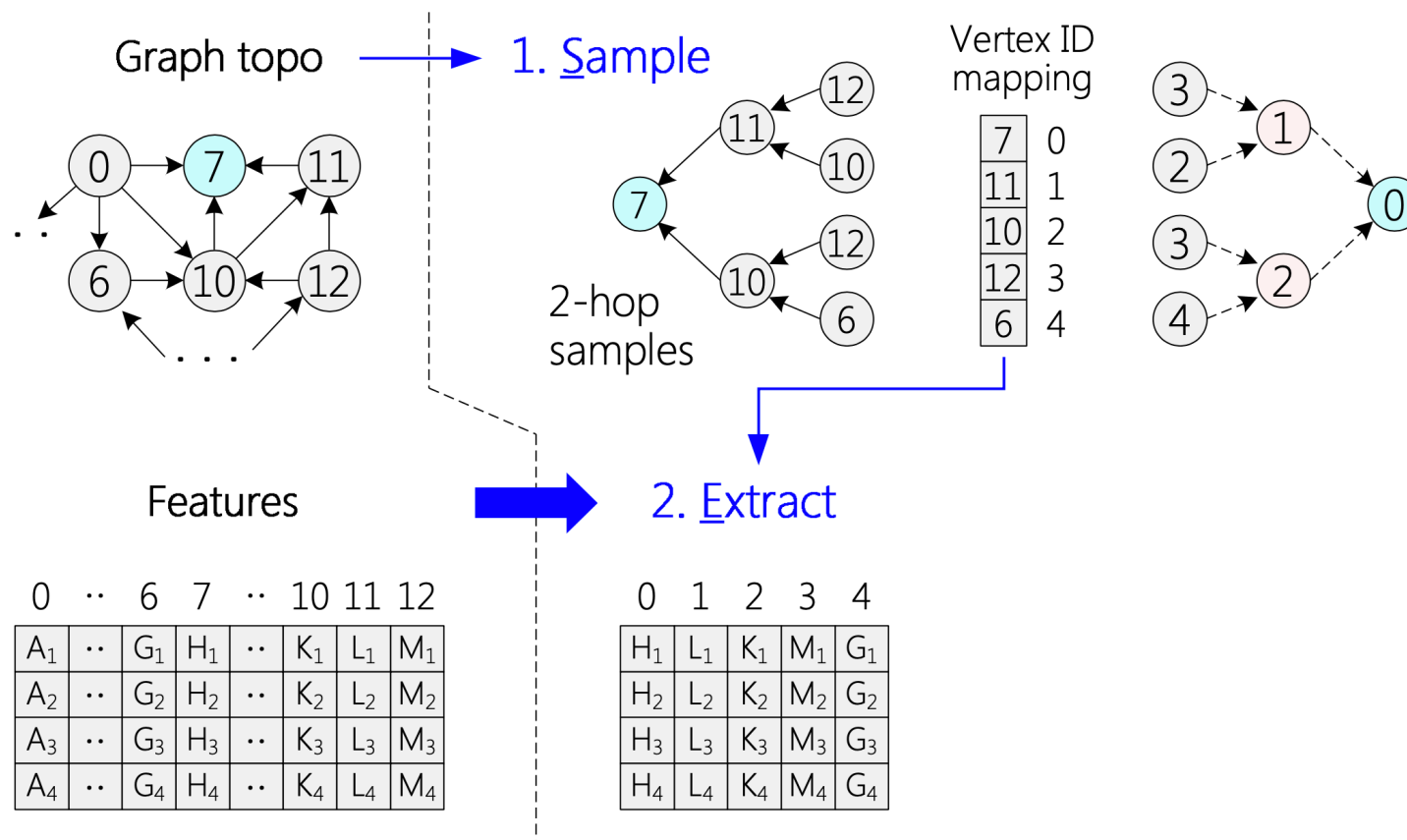
Sample-base GNN Training



7

SET model

1. Sample
2. Extract



E

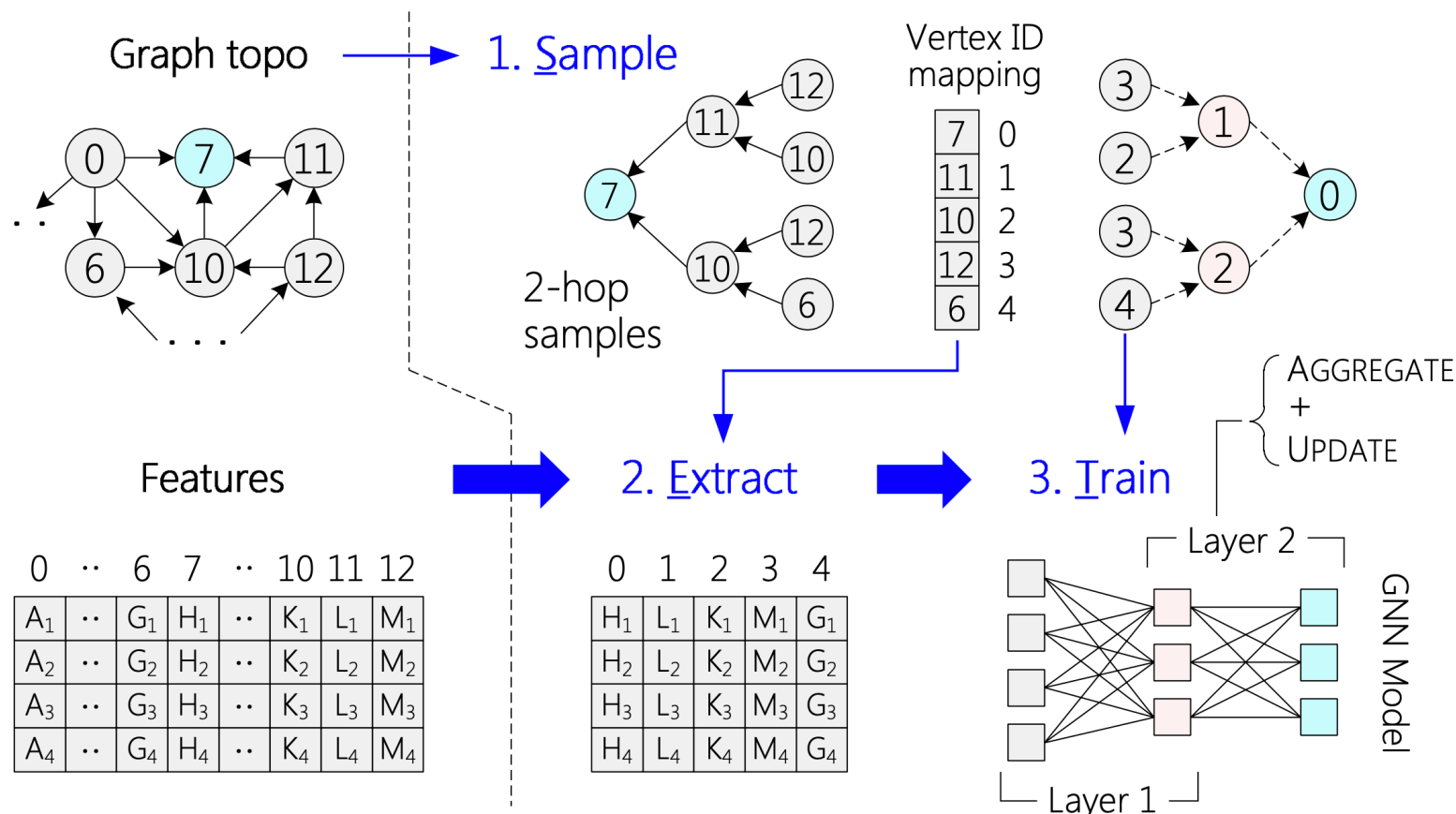
Sample-base GNN Training



8

SET model

1. Sample
2. Extract
3. Train



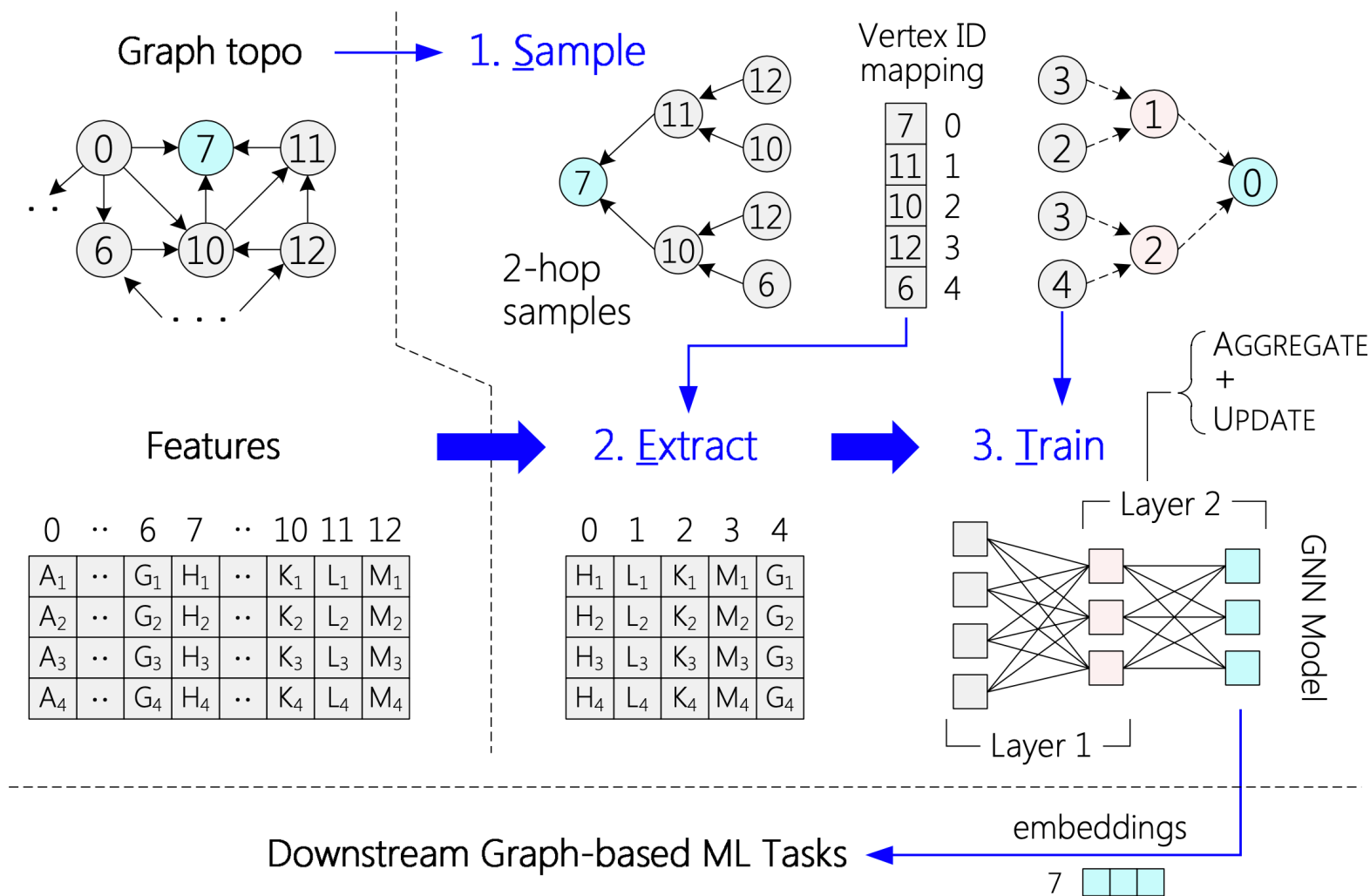
Sample-base GNN Training



9

SET model

1. Sample
2. Extract
3. Irain



GPU-based GNN Training



10

GPUs have been widely exploited to accelerate GNN training

► Train: almost

GNN Systems	Sample	Extract	Train	TOT
T_{SOTA}	2.93	5.55	4.00	12.50

GCN w/ 3-hop sampling on OGB-Papers100M

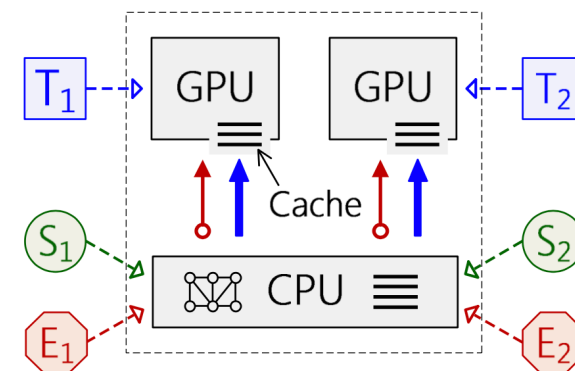
GPU-based GNN Training



11

GPUs have been widely exploited to accelerate GNN training

- ▶ Train: almost
- ▶ Extract: PaGraph^[SOCC'20]



GNN Systems	Sample	Extract	Train	TOT
T _{SOTA}	2.93	5.55	4.00	12.50
w/ GPU-based Caching	2.88	1.73	4.00	8.62

GCN w/ 3-hop sampling on OGB-Papers100M

GPU-based GNN Training

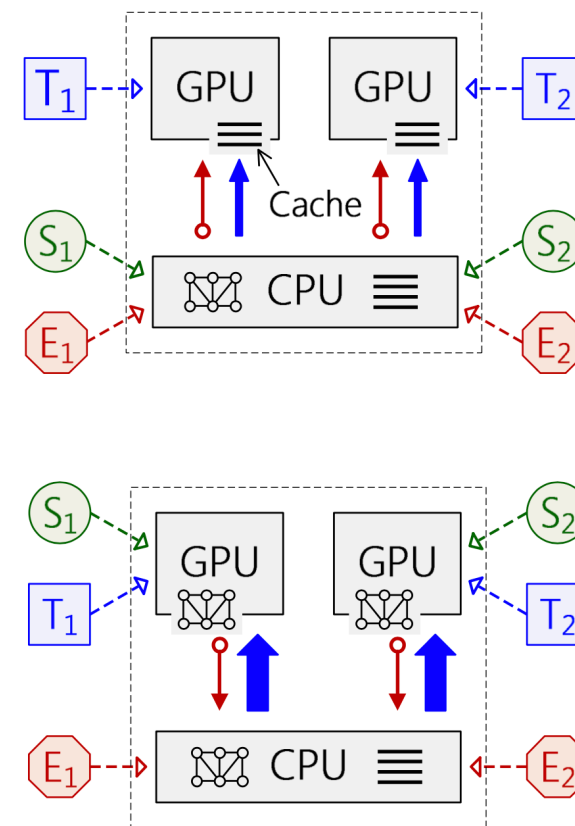
12

GPUs have been widely exploited to accelerate GNN training

- ▶ Train: almost
- ▶ Extract: PaGraph^[SOCC'20]
- ▶ Sample: DGL, NextDoor^[EuroSys'21]

GNN Systems	Sample	Extract	Train	TOT
T_{SOTA}	2.93	5.55	4.00	12.50
w/ GPU-based Caching	2.88	1.73	4.00	8.62
w/ GPU-based Sampling	0.70	5.46	4.01	10.21

GCN w/ 3-hop sampling on OGB-Papers100M



GPU-based GNN Training

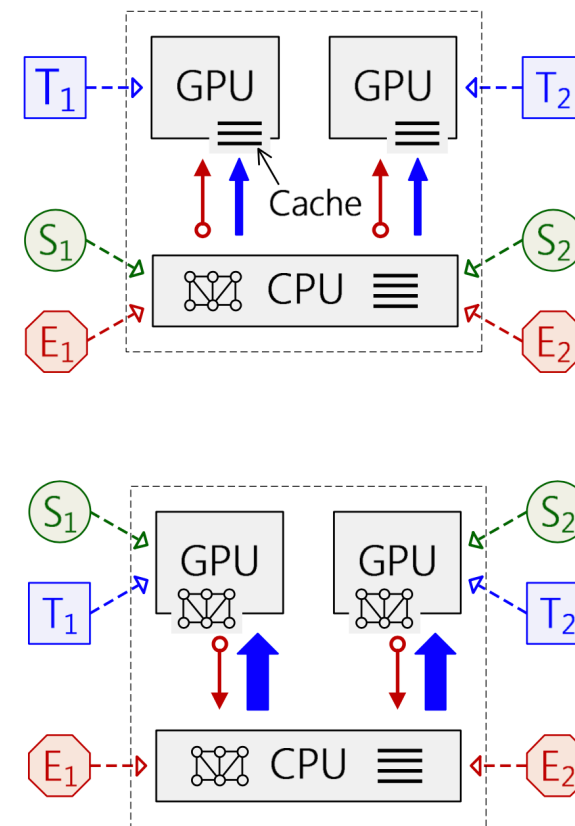
13

GPUs have been widely exploited to accelerate GNN training

- ▶ Train: almost
- ▶ Extract: PaGraph^[SOCC'20]
- ▶ Sample: DGL, NextDoor^[EuroSys'21]
- ▶ **Both ?**

GNN Systems	Sample	Extract	Train	TOT
T_{SOTA}	2.93	5.55	4.00	12.50
w/ GPU-based Caching	2.88	1.73	4.00	8.62
w/ GPU-based Sampling	0.70	5.46	4.01	10.21
w/ Both	0.70	3.62	4.00	8.37

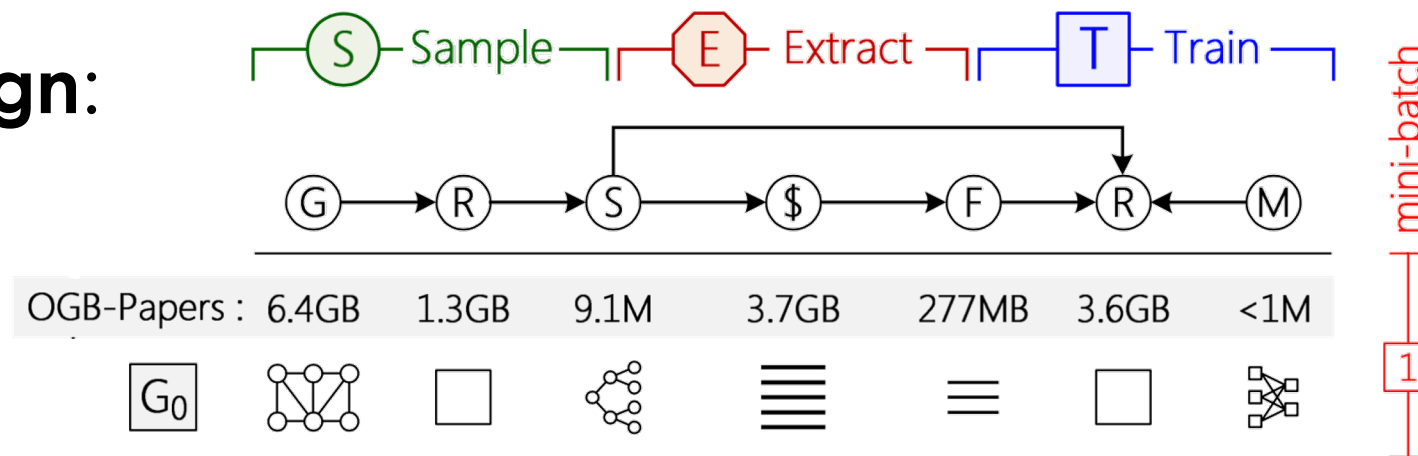
GCN w/ 3-hop sampling on OGB-Papers100M



Analysis

14

Traditional Design:

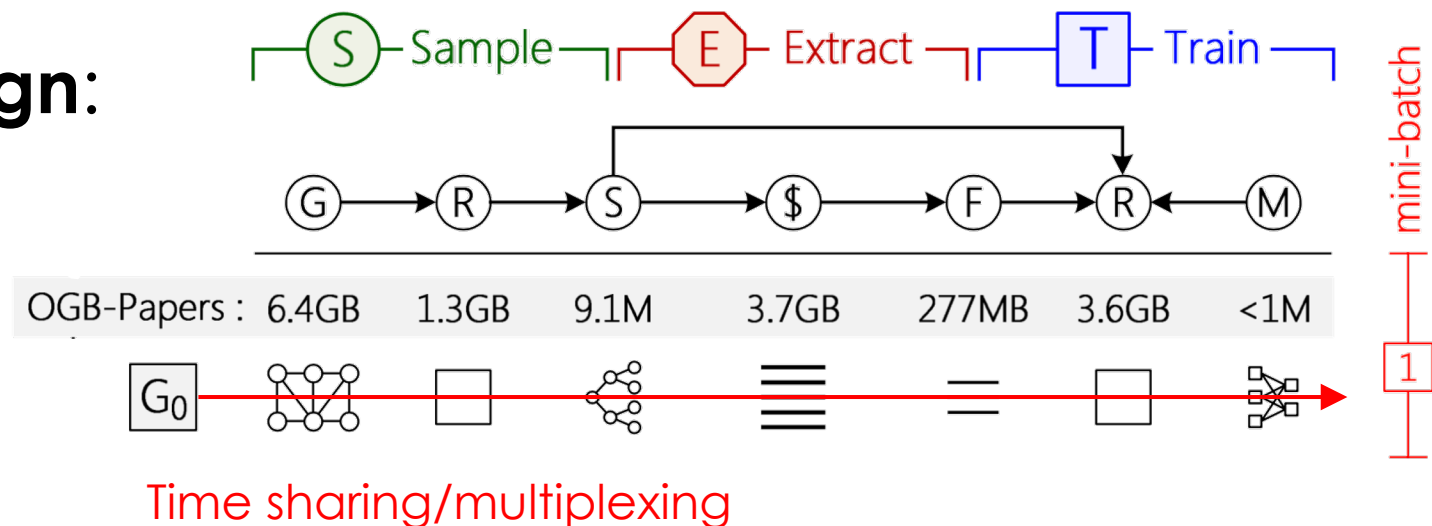


(G) Graph topo (S) Samples (\$) Cache (F) Features (M) Model (R) Runtime state

Analysis

15

Traditional Design:



(G) Graph topo (S) Samples (\$) Cache (F) Features (M) Model (R) Runtime state

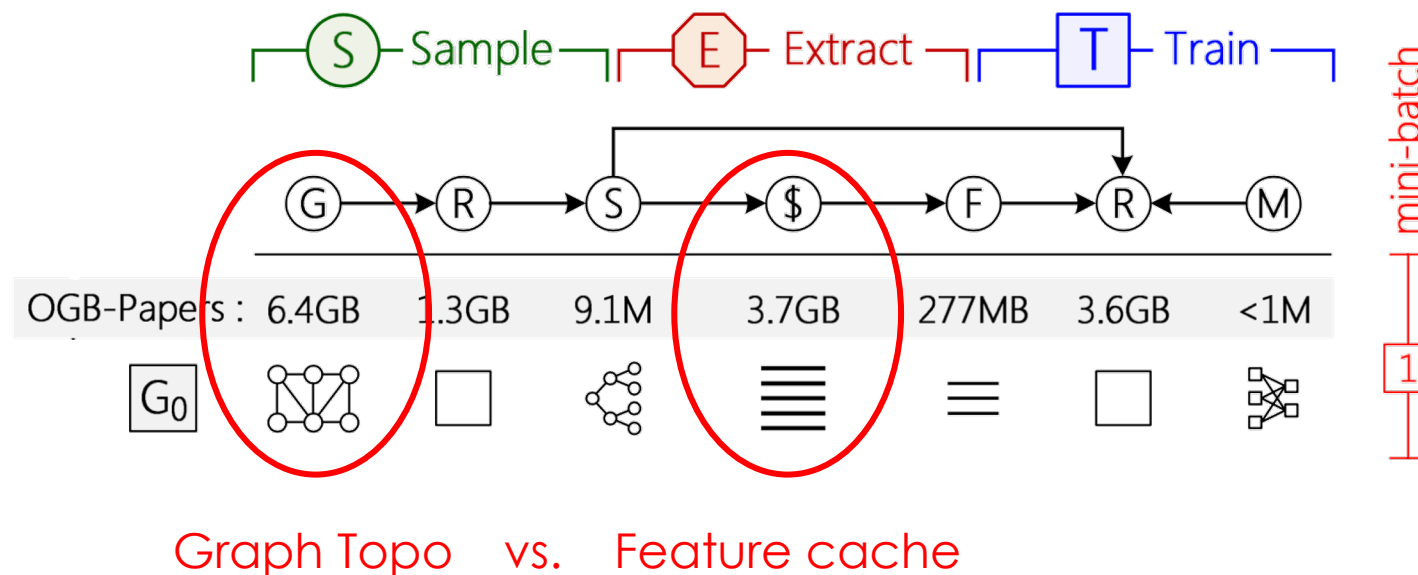
Analysis



16

Problems

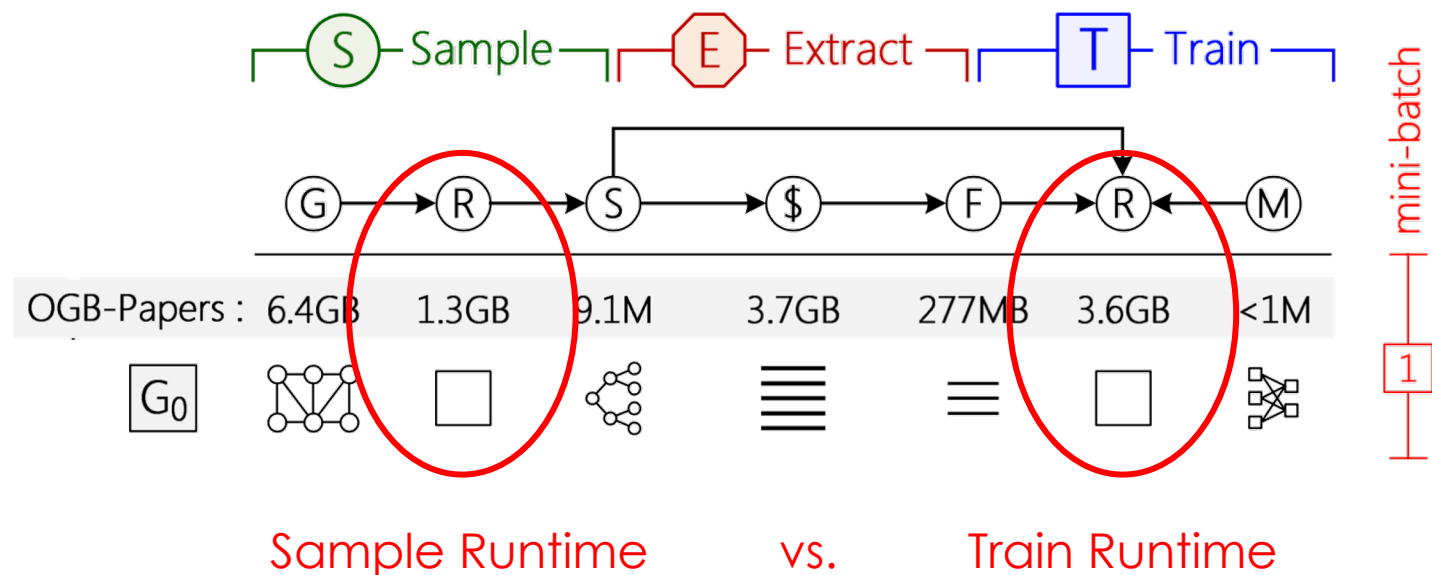
► Capacity



Analysis

Problems

► Capacity



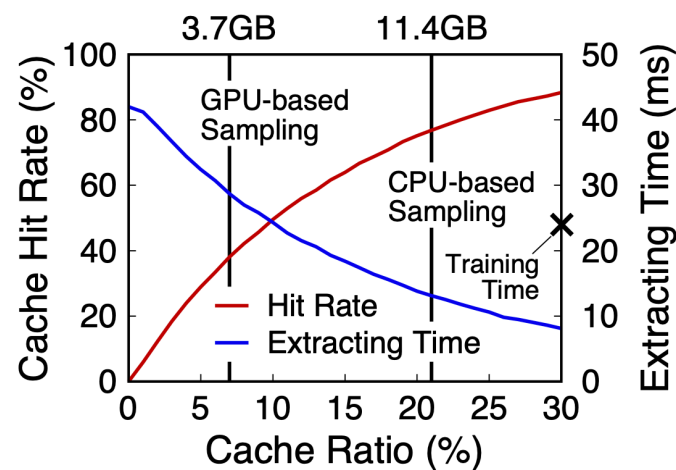
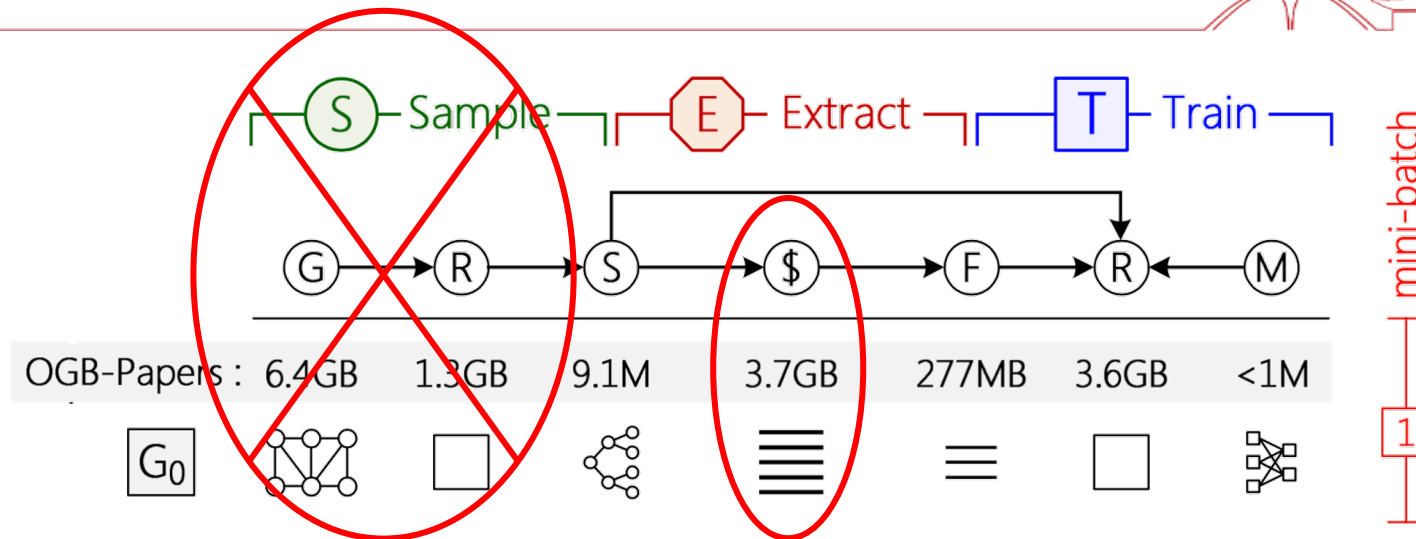
① Graph topo ② Samples ③ Cache ④ Features ⑤ Model ⑥ Runtime state

Analysis

18

Problems

► Capacity



OGB-Papers100M

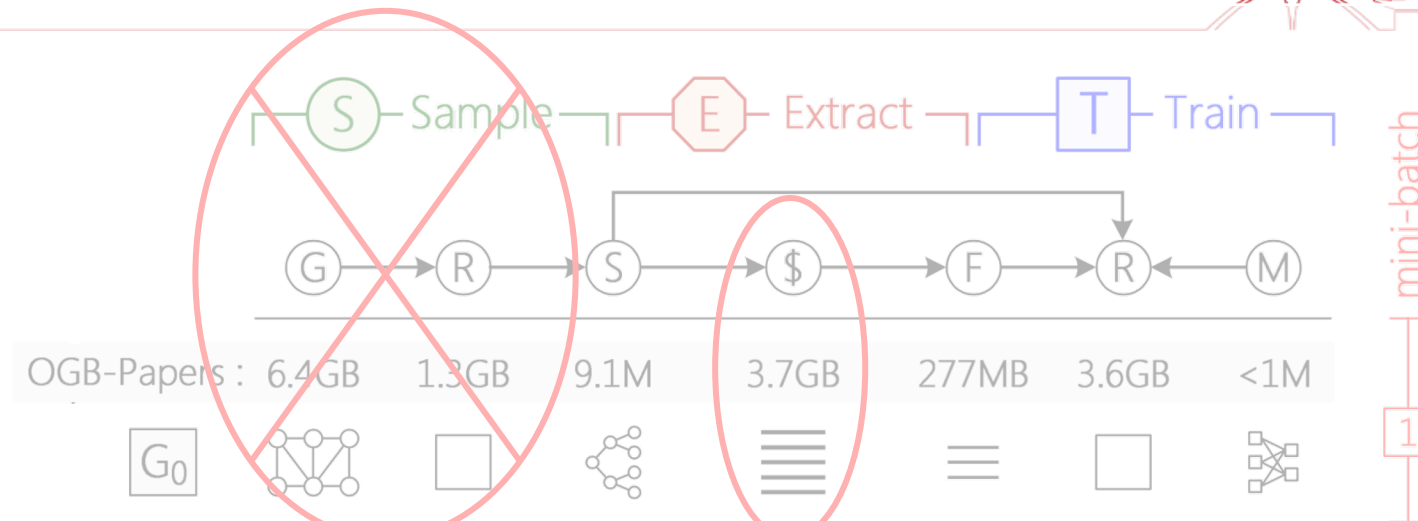
(G) Graph topo (S) Samples (\$) Cache (F) Features (M) Model (R) Runtime state

Analysis

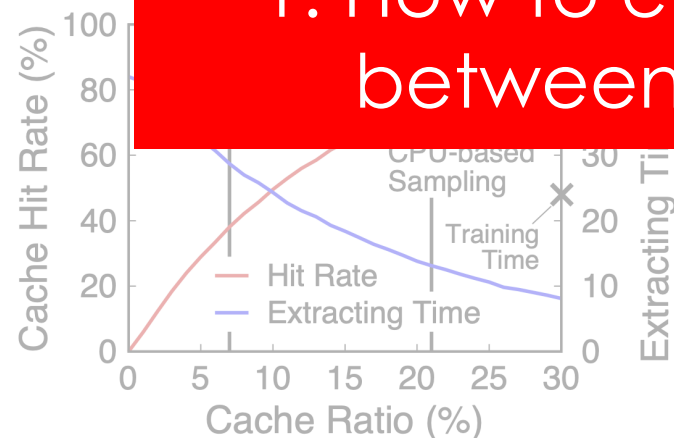
19

Problems

► Capacity



1. How to eliminate **contention on GPU memory** between different stages of the SET model



OGB-Papers100M

11.4 GB

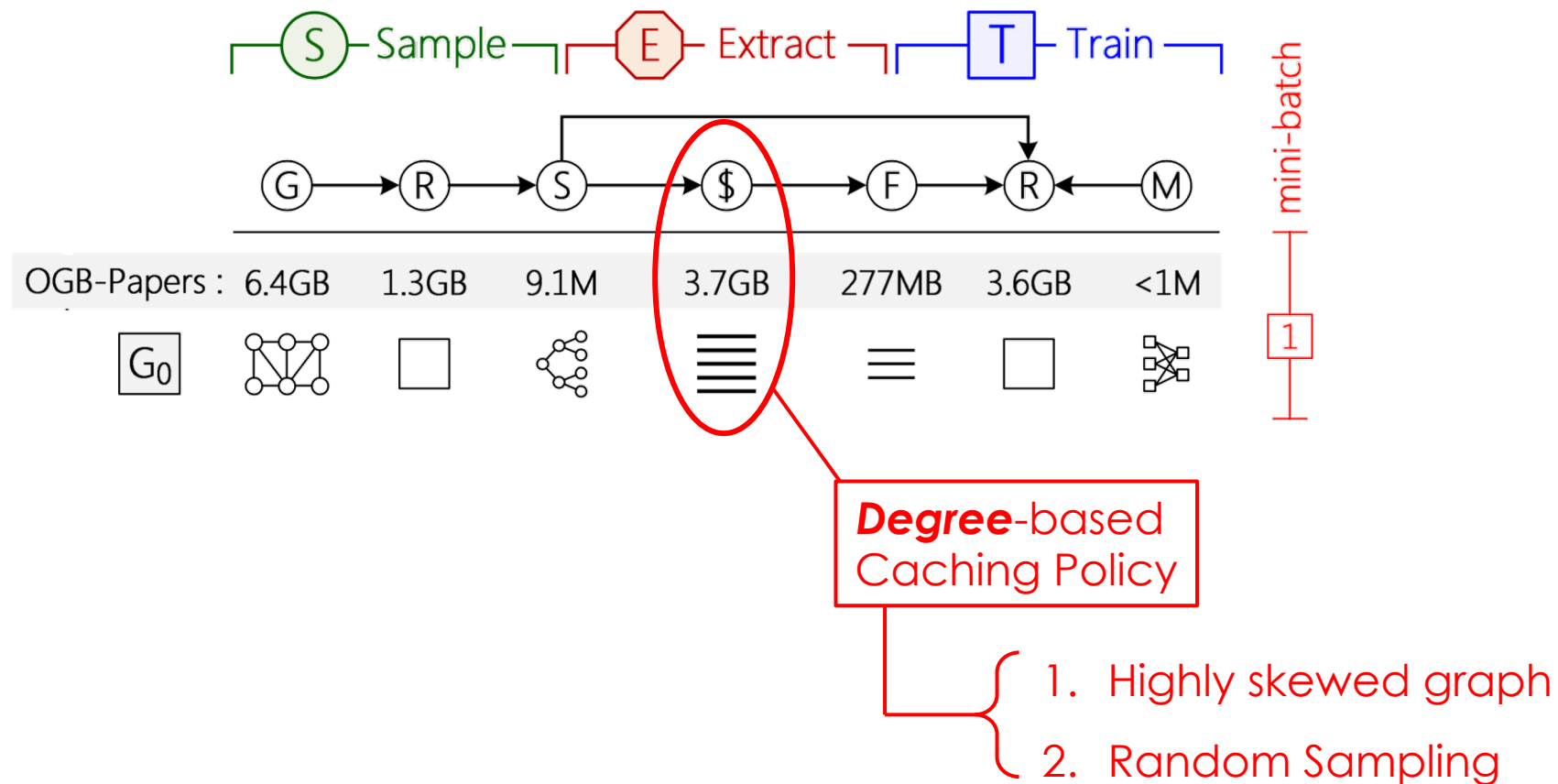
(G) Graph topo (S) Samples (\$) Cache (F) Features (M) Model (R) Runtime state

Analysis

20

Problems

- Capacity
- Efficiency

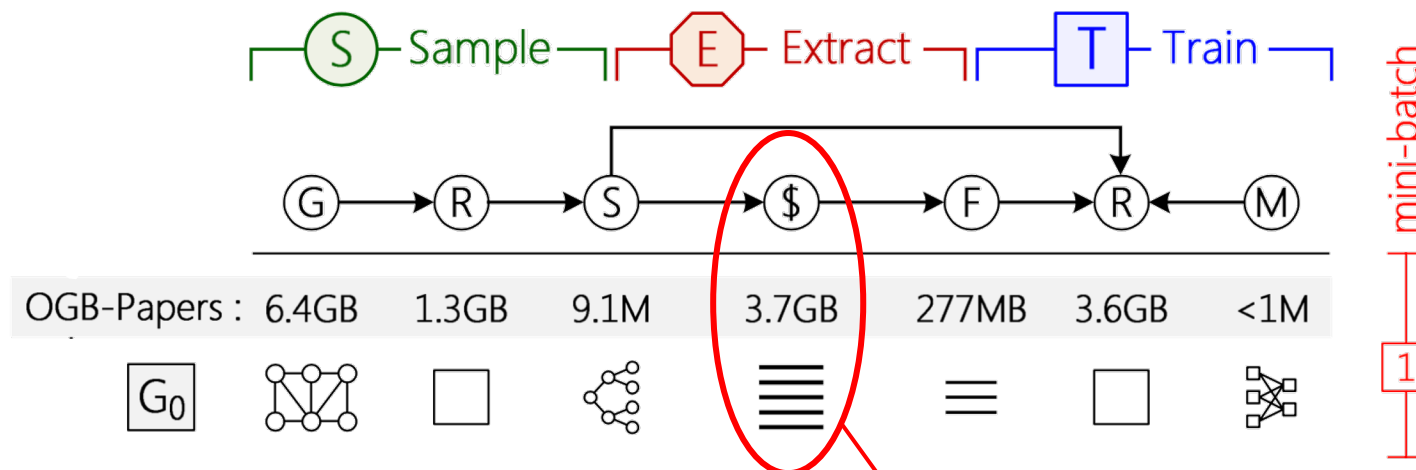


Analysis

21

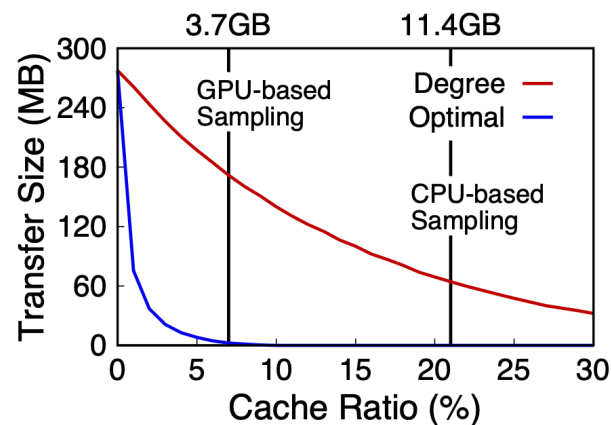
Problems

- Capacity
- Efficiency

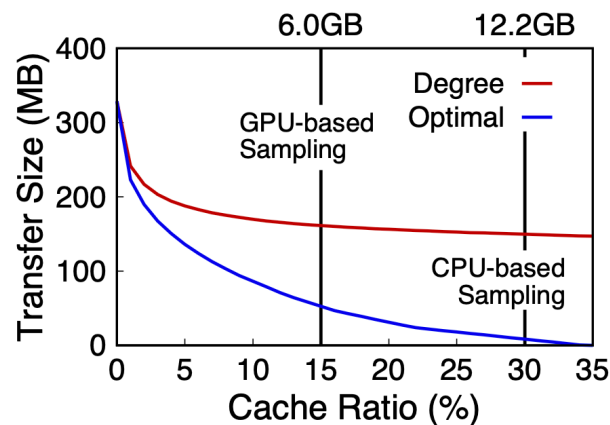


Degree-based Caching Policy

- 1. Highly skewed graph
- 2. Random Sampling



OGB-Papers100M



Weighted sampling

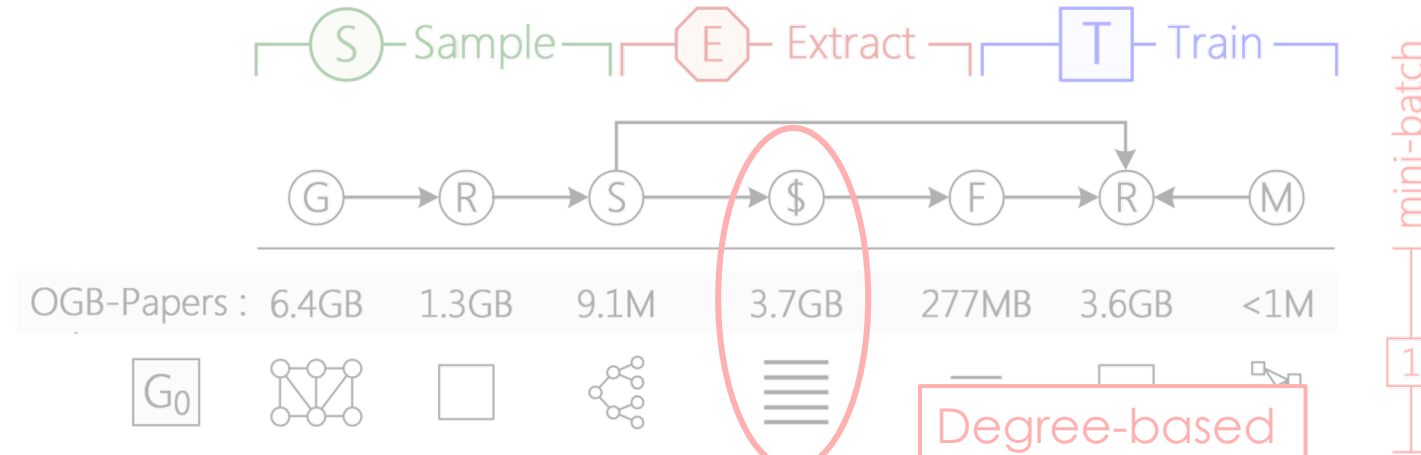
(G) Graph topo (S) Samples (\$) Cache (F) Features (M) Model (R) Runtime state

Analysis

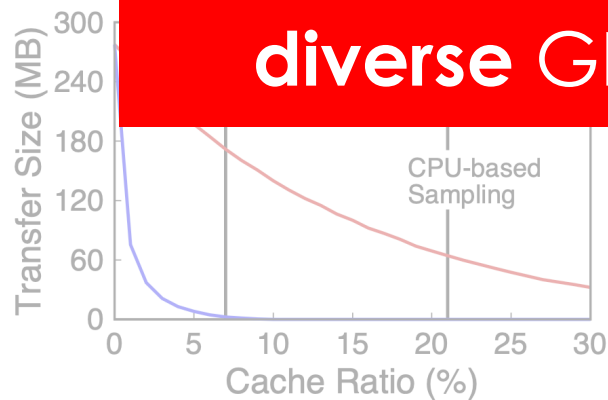
22

Problems

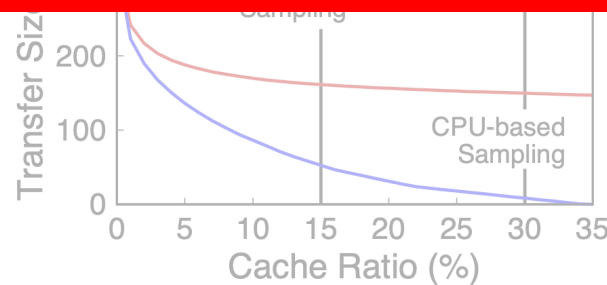
- Capacity
- Efficiency



2. How to achieve **optimal cache efficiency** for **diverse** GNN datasets and sampling algorithms



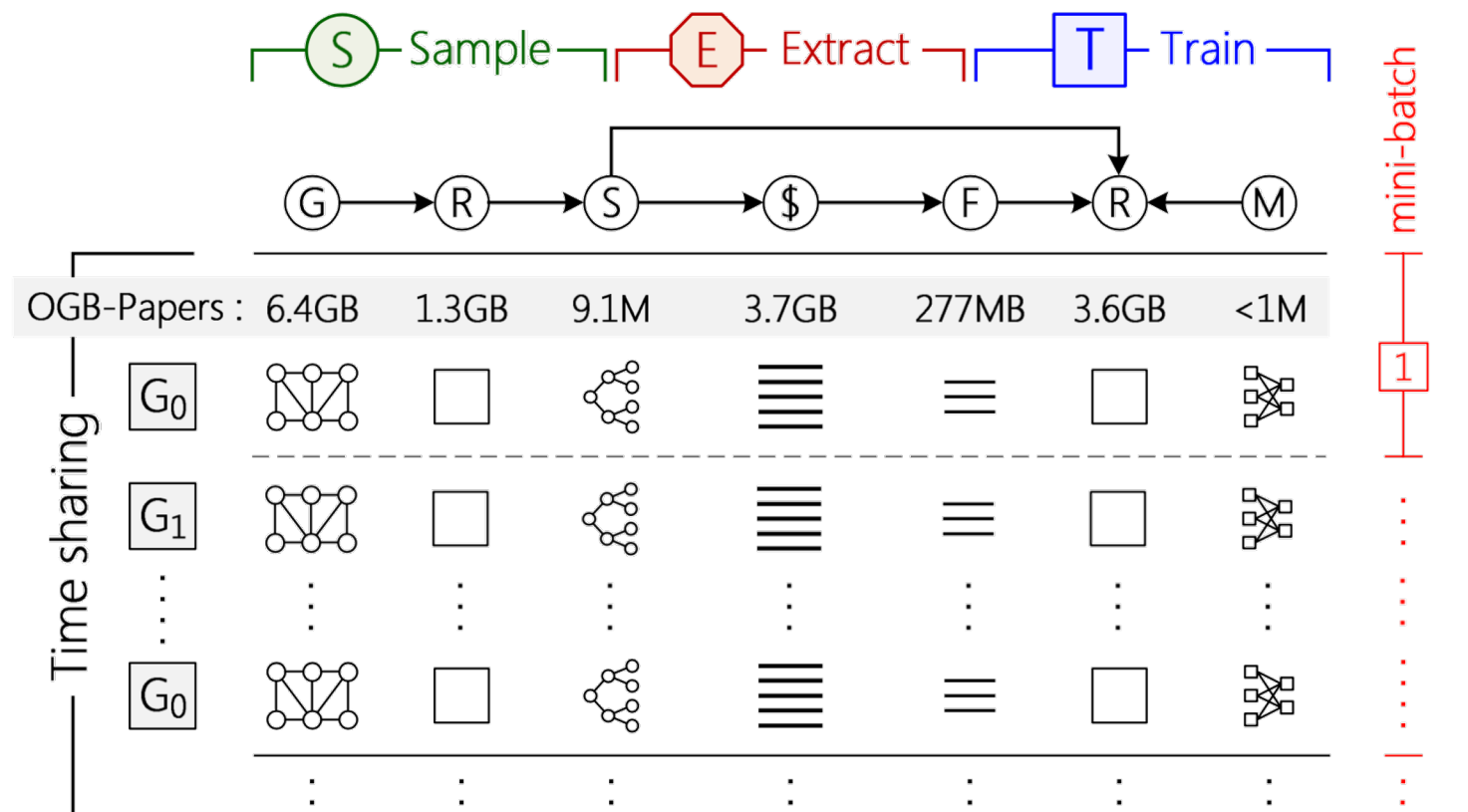
OGB-Papers100M



Weighted sampling

General Idea

Time sharing w/ multi-GPUs



Intra-task Contention



General Idea

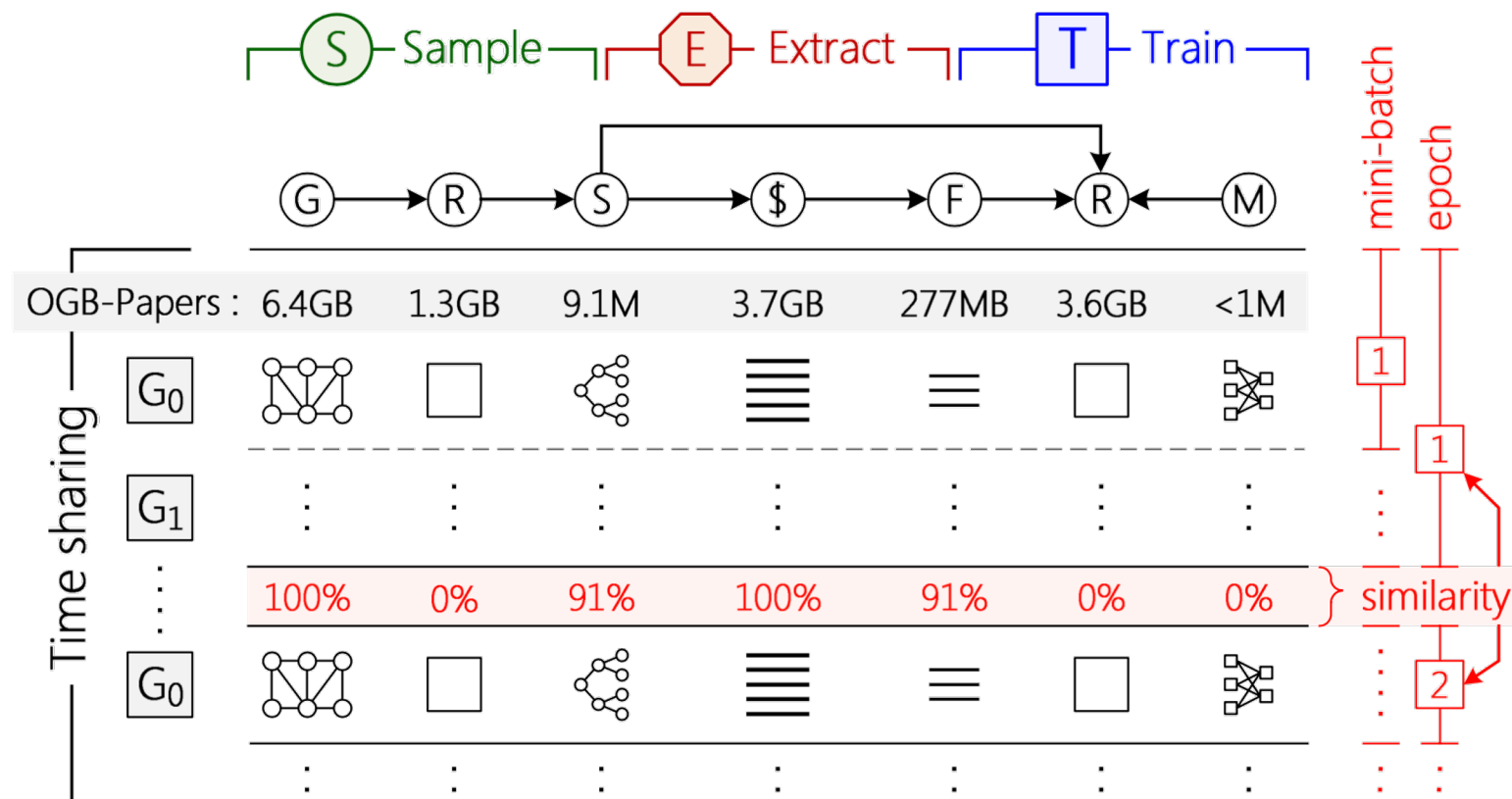
25

Observation:

HIGH

cross-GPU

similarity



(G) Graph topo (S) Samples (\$) Cache (F) Features (M) Model (R) Runtime state

General Idea

26

Observation:

HIGH

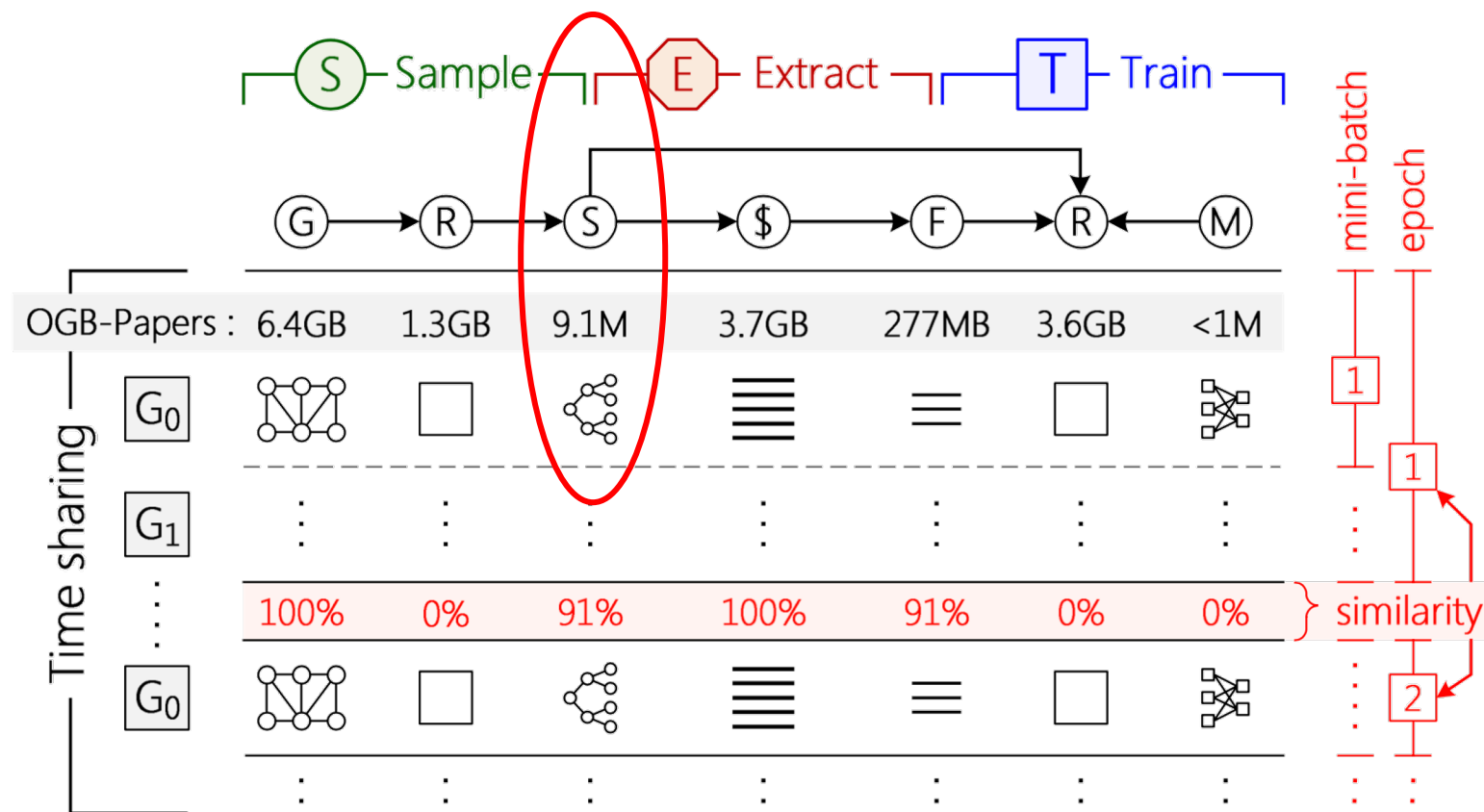
cross-GPU

similarity

LOW

cross-stage

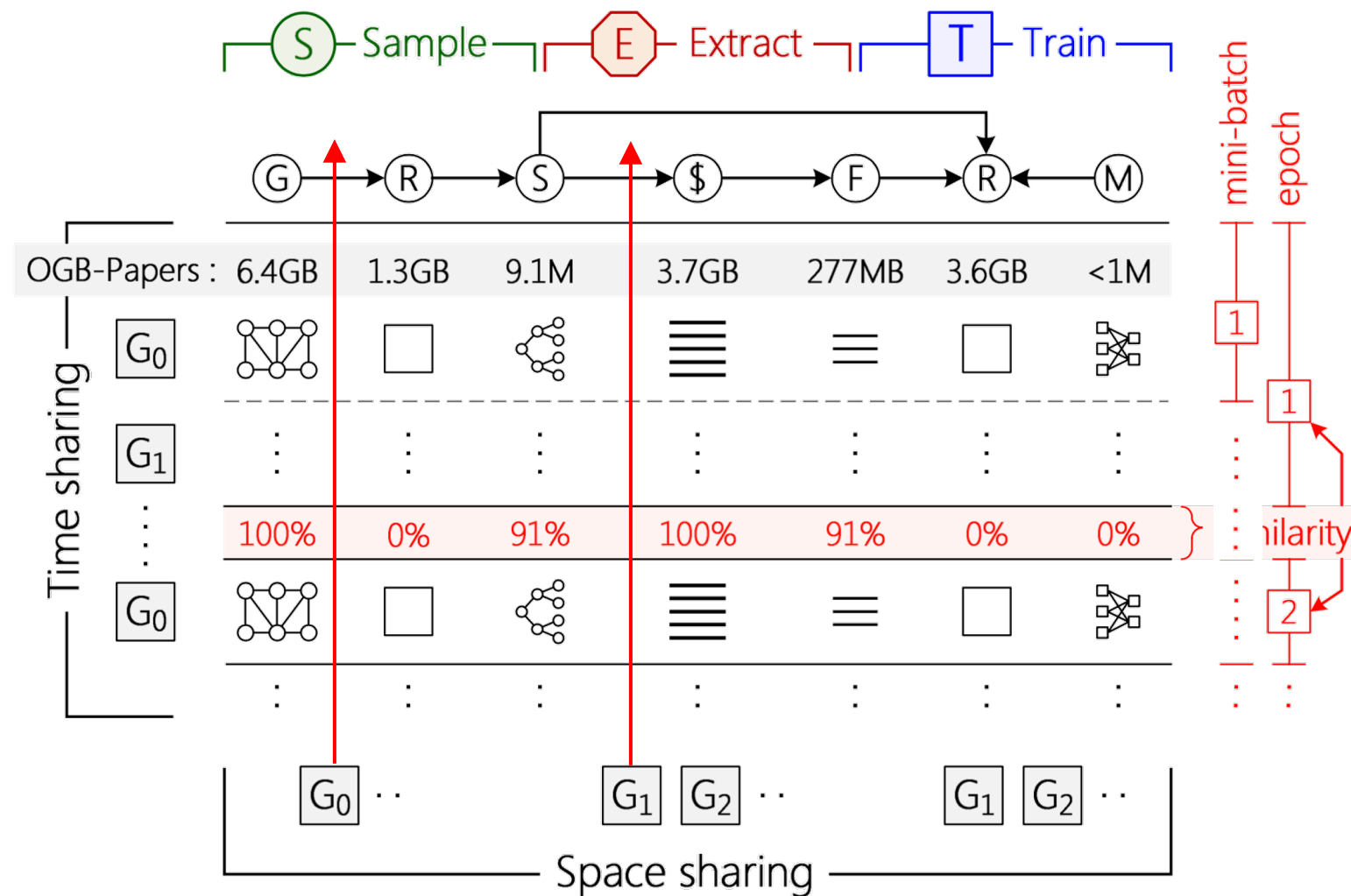
data sharing



(G) Graph topo (S) Samples (\$) Cache (F) Features (M) Model (R) Runtime state

General Idea

Space sharing w/ multi-GPUs



David Wentzlaff and Anant Agarwal @MIT

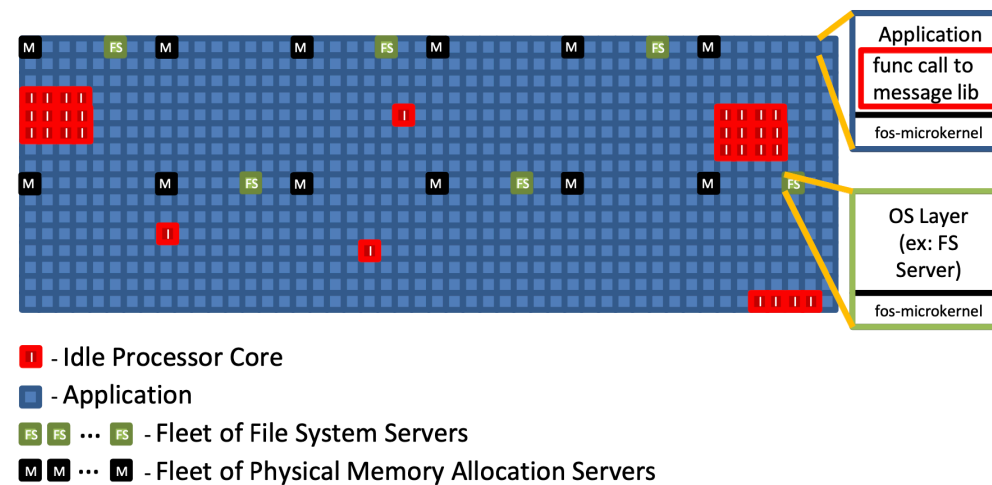


Figure 4. OS and application clients executing on the fos-microkernel

Our Approach

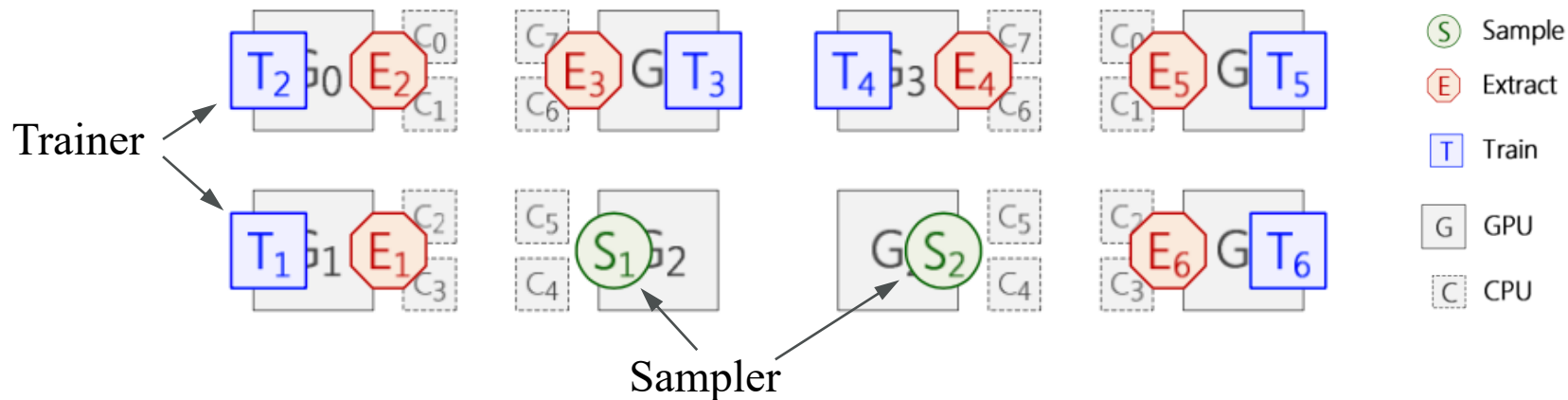


29

A factored design

- ▶ Inspired by the **factored operating system** ($fos^{[ACM\ OSR'09]}$)
- ▶ **GNNLab**: a factored system for sample-based GNN training
 - ❖ Perform each stage on **dedicate processors** (GPUs and/or CPUs)

An example of GNNLab on an 8-GPU machine (2 Samplers & 6 Trainers)



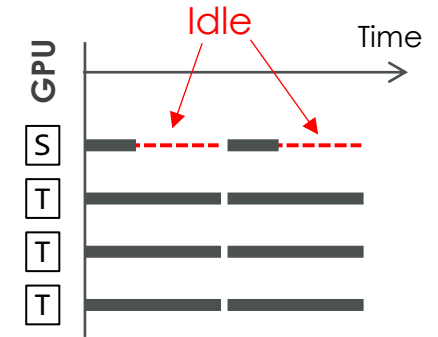
Our Approach



30

Fundamental Challenge: **load imbalance**

- ▶ **Coarse-grained** (stage-level) workload partitioning
- ▶ **Limited GPUs**: normally ≤ 8 , even just 1
- ▶ **Diverse** datasets and workloads
 - ❖ Sampling vs. Training: e.g., GCN = 1 : 1, while PinSAGE = 1 : 10



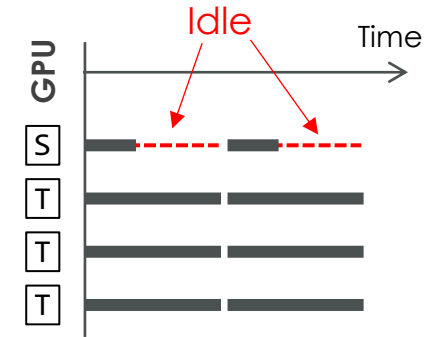
Our Approach



31

Fundamental Challenge: **load imbalance**

- ▶ **Coarse-grained** (stage-level) workload partitioning
- ▶ **Limited GPUs**: normally ≤ 8 , even just 1
- ▶ **Diverse** datasets and workloads
 - ❖ Sampling vs. Training: e.g., GCN = 1 : 1, while PinSAGE = 1 : 10



GOALs of GNNLab

1. Make stages work together **efficiently**
2. Assign GPUs to different stages **flexibly**

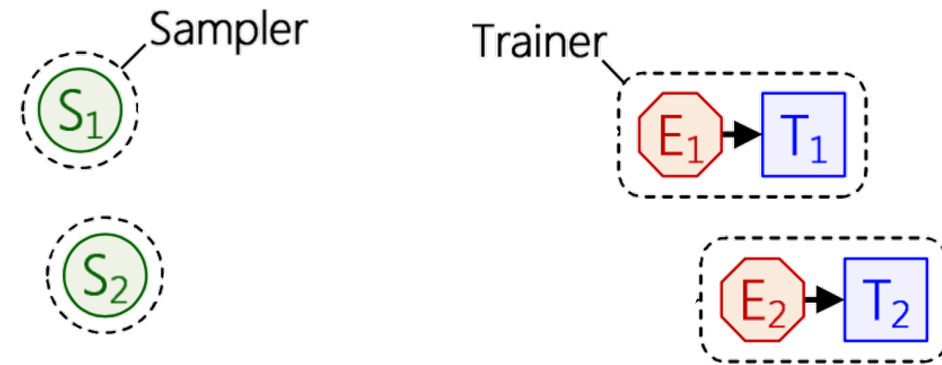
Execution Engine



32

Architecture

- ▶ Two executors **@GPUs**
 - ❖ **Sampler**: Sample stage
 - ❖ **Trainer**: Extract & Train stage

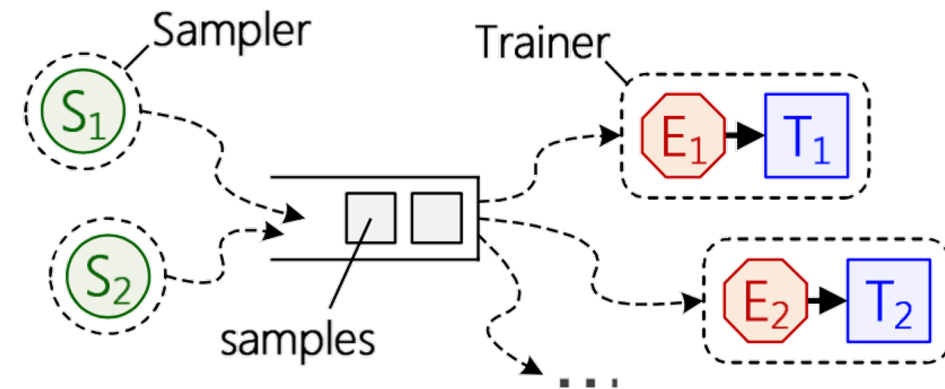


Execution Engine



Architecture

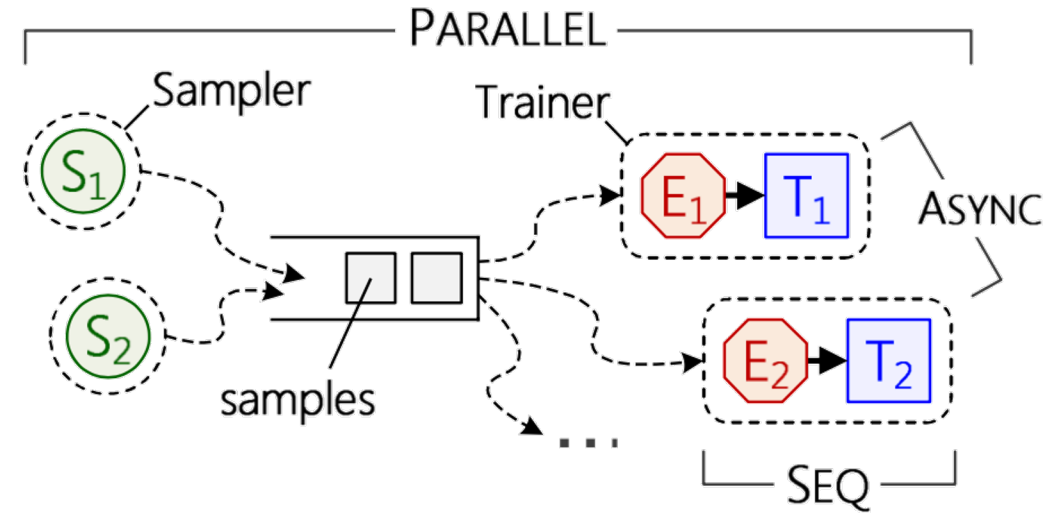
- ▶ Two executors @GPUs
 - ❖ **Sampler**: Sample stage
 - ❖ **Trainer**: Extract & Train stage
- ▶ A **global** queue @CPUs



Execution Engine

Architecture

- ▶ Two executors @GPUs
 - ❖ **Sampler**: Sample stage
 - ❖ **Trainer**: Extract & Train stage
- ▶ A **global** queue @CPUs
- ▶ Execution flow
 - ❖ Inter-executor: **Parallel**
 - ❖ Intra-executor: **Sequential** w/ **pipelining**
 - ❖ Gradient updates w/ bounded staleness



GPU allocation scheme

- ▶ OB: performance of executors on GPU is **predictable**
- ▶ N_g : the number of GPUs
- ▶ N_s (resp. N_t): #GPUs allocated to Samplers (resp. Trainers)
- ▶ T_s (resp. T_t): the processing time of Sampler (resp. Trainer)

$$N_s = \left\lceil \frac{N_g}{K + 1} \right\rceil$$

$$K = \frac{T_t}{T_s}$$

GPU allocation scheme

- ▶ OB: performance of executors on GPU is **predictable**
- ▶ N_g : the number of GPUs
- ▶ N_s (resp. N_t): #GPUs allocated to Samplers (resp. Trainers)
- ▶ T_s (resp. T_t): the processing time of Sampler (resp. Trainer)
- ▶ **Prefer to** allocate GPUs to Samplers
 - ❖ *temporarily switching from Sampler to Trainer is efficient*
- ▶ Dynamic executor **switching** [see our paper]

Prefer to Sampler

$$N_s = \left\lfloor \frac{N_g}{K + 1} \right\rfloor$$

$$K = \frac{T_t}{T_s}$$

A general caching scheme

- ▶ Hotness metric h_v and cache ratio α
 1. Store and sort vertices according to their h_v
 2. Load features of *top-ranked* $\alpha |V|$ vertices w.r.t. h_v into GPU cache

For example: PaGraph^[SOCC'20]

h_v = out-degree of each vertex

Pre-sampling based caching policy (*PreSC*)

- OB: Cross-task (epoch) *similarity* of access footprint

Sampling alogrithms	PR	TW	PA	UK
3-hop random	73.97	78.89	91.29	77.46
Random walks	78.16	72.68	87.14	64.40
3-hop weighted	77.69	66.64	89.57	72.96

Pre-sampling based caching policy (*PreSC*)

- ▶ OB: Cross-task (epoch) **similarity** of access footprint

Sampling algorithms	PR	TW	PA	UK
3-hop random	73.97	78.89	91.29	77.46
Random walks	78.16	72.68	87.14	64.40
3-hop weighted	77.69	66.64	89.57	72.96

- ▶ General idea: **pre-sample** a few rounds to estimate vertex **hotness**
 1. Conducts **K** sampling stages (normally 1) for the training set
 2. Record **visit count** of the sampled vertices
 3. Use average count as the hotness metric **h_v**

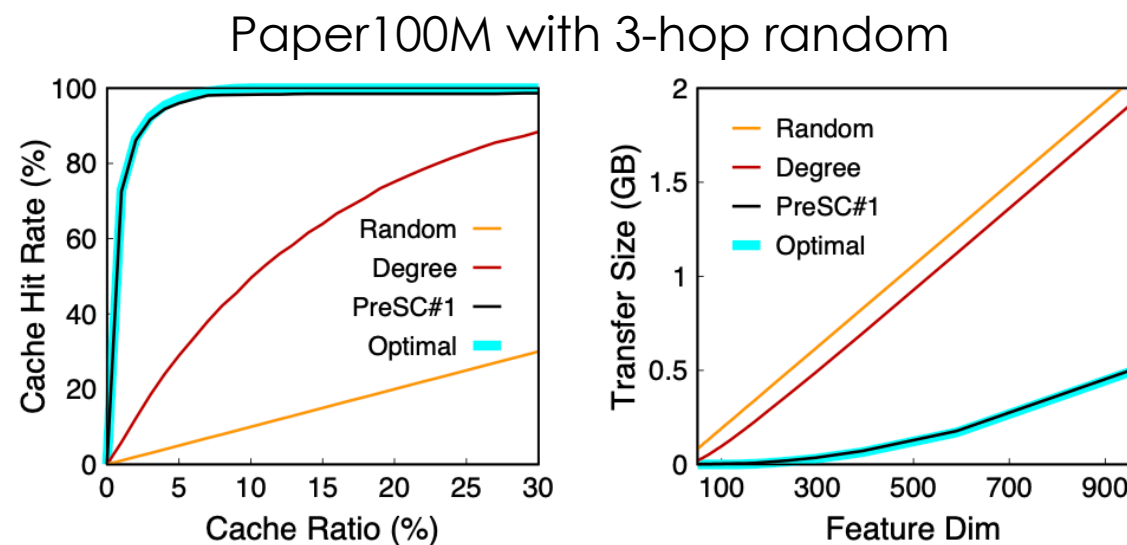
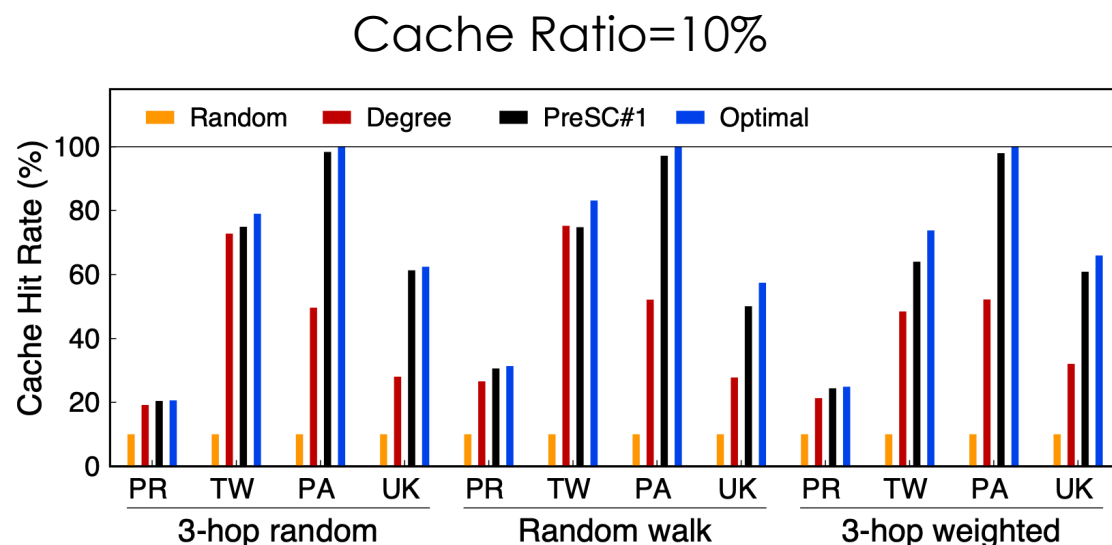
GPU-based Feature Caching



40

Pre-sampling based caching policy (*PreSC*)

- **Efficiency**: close to Optimal, vs. Degree avg $1.5\times$ (up to $2.2\times$)
- **Robustness**: stable for all 12 cases



Testbed

- ▶ 8 NVIDIA V100 GPU
w/ 16GB memory
- ▶ Intel Xeon 2×24 CPU

GNNs

- ▶ GCN (3-hop rand ngb)
- ▶ GraphSAGE (2-hop rand ngb)
- ▶ PinSAGE (rand walks)

Datasets

Dataset	#Vertex	#Edge	Dim.	#TS	Vol _G	Vol _F
PR [5]	2.4M	124M	100	197K	481MB	934MB
TW [33]	41.7M	1.5B	256	417K	5.6GB	40GB
PA [4]	111M	1.6B	128	1.2M	6.4GB	53GB
UK [9]	77.7M	3.0B	256	1.0M	11.3GB	74GB

PR can be loaded into a single GPU

Baselines

System	Design	Sample	Extract	Train
PyG	N/A	CPU	No cache	GPU
DGL	Time S.	GPU	No cache	GPU
T _{SOTA}	Time S.	GPU w/ Opt.	Cache w/ Degree	GPU
GNNLab	Space S.	GPU w/ Opt.	Cache w/ PreSC	GPU

Overall Performance



42

► MERITS:

(A1) space sharing design

(A2) pre-sampling policy

(A3) efficient sampling impl.

► vs. PyG: $10.2\times \sim 74.3\times$

► vs. DGL: $2.4\times \sim 9.1\times$
due to (A1)~(A3)

► vs. T_{SOTA} : $1.6\times \sim 3.8\times$, e.f. PR
due to (A1) and (A2)

Our flexible scheduling scheme already provides
optimal GPU allocations in an 8-GPU machine

GNN Model	Dataset	PyG	DGL	T_{SOTA}	GNNLab
GCN	PR	11.91	1.33	0.22	0.33 (2S)
	TW	12.15	3.86	1.80	0.47 (2S)
	PA	14.82	4.56	2.46	0.84 (2S)
	UK	15.04	OOM	OOM	1.47 (2S)
GraphSAGE	PR	8.17	0.79	0.07	0.11 (4S)
	TW	8.18	1.81	0.35	0.20 (2S)
	PA	9.68	2.47	0.85	0.30 (2S)
	UK	9.86	OOM	2.01	0.61 (1S)
PinSAGE	PR	$\times$	0.94	0.30	0.40 (1S)
	TW	$\times$	2.50	0.98	0.58 (1S)
	PA	$\times$	2.97	1.65	1.05 (1S)
	UK	$\times$	OOM	OOM	1.81 (1S)

Performance Breakdown



43

S, **E**, and **T** represent Sample, Extract, and Train stages. **G**, **M**, and **C** represent graph sampling, marking cached vertices, and copying samples to host memory in Sample stage, respectively. **R%** and **H%** represent the cache ratio of features and the cache hit rate.

GNN	Dataset	DGL			T _{SOTA}			GNNLab		
		<u>S</u>	<u>E</u>	<u>T</u>	<u>S</u> = G + M	<u>E</u> (R%, H%)	<u>T</u>	<u>S</u> = G + M + C	<u>E</u> (R%, H%)	<u>T</u>
GCN	PR	0.35	2.81	1.22	0.30 = 0.29 + 0.01	0.04 (100, 100)	1.18	0.39 = 0.29 + 0.01 + 0.09	0.15 (100, 100)	1.18
	TW	0.74	9.44	1.48	0.29 = 0.26 + 0.03	3.68 (1, 29)	1.53	0.37 = 0.26 + 0.03 + 0.08	0.76 (25, 89)	1.51
	PA	1.20	10.70	4.00	0.79 = 0.70 + 0.10	3.64 (7, 38)	4.00	0.96 = 0.68 + 0.10 + 0.18	0.49 (21, 99)	3.82
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.56 = 0.39 + 0.03 + 0.14	3.06 (14, 70)	3.09
GSG	PR	0.13	1.92	0.23	0.16 = 0.15 + 0.01	0.03 (100, 100)	0.25	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	0.24
	TW	0.38	4.65	0.44	0.12 = 0.11 + 0.01	0.62 (15, 77)	0.44	0.16 = 0.11 + 0.01 + 0.03	0.41 (32, 89)	0.43
	PA	0.56	6.06	1.25	0.38 = 0.33 + 0.06	1.42 (11, 56)	1.18	0.46 = 0.31 + 0.06 + 0.08	0.28 (25, 99)	1.15
	UK	OOM	OOM	OOM	0.19 = 0.19 + 0.00	4.49 (0, 0)	1.08	0.26 = 0.18 + 0.02 + 0.06	1.39 (18, 72)	1.01
PSG	PR	0.40	1.64	1.75	0.16 = 0.16 + 0.01	0.03 (100, 100)	1.74	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	1.72
	TW	0.72	5.22	2.59	0.23 = 0.22 + 0.02	1.12 (4, 60)	2.60	0.28 = 0.21 + 0.02 + 0.05	0.51 (26, 86)	2.52
	PA	1.86	4.85	5.78	0.54 = 0.49 + 0.05	1.68 (6, 37)	6.09	0.61 = 0.47 + 0.04 + 0.09	0.33 (22, 97)	6.01
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.65 = 0.49 + 0.03 + 0.13	3.37 (13, 57)	7.00

Performance Breakdown



44

S, **E**, and **T** represent Sample, Extract, and Train stages. **G**, **M**, and **C** represent graph sampling, marking cached vertices, and copying samples to host memory in Sample stage, respectively. **R%** and **H%** represent the cache ratio of features and the cache hit rate.

GNN	Dataset	DGL			T _{SOTA}			GNNLab		
		<u>S</u>	<u>E</u>	<u>T</u>	<u>S</u> = G + M	<u>E</u> (R%, H%)	<u>T</u>	<u>S</u> = G + M + C	<u>E</u> (R%, H%)	<u>T</u>
GCN	PR	0.35	2.81	1.22	0.30 = 0.29 + 0.01	0.04 (100, 100)	1.18	0.39 = 0.29 + 0.01 + 0.09	0.15 (100, 100)	1.18
	TW	0.74	9.44	1.48	0.29 = 0.26 + 0.03	3.68 (1, 29)	1.53	0.37 = 0.26 + 0.03 + 0.08	0.76 (25, 89)	1.51
	PA	1.20	10.70	4.00	0.79 = 0.70 + 0.10	3.64 (7, 38)	4.00	0.96 = 0.68 + 0.10 + 0.18	0.49 (21, 99)	3.82
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.56 = 0.39 + 0.03 + 0.14	3.06 (14, 70)	3.09
GSG	PR	0.13	1.92	0.23	0.16 = 0.15 + 0.01	0.03 (100, 100)	0.25	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	0.24
	TW	0.38	4.65	0.44	0.12 = 0.11 + 0.01	0.62 (15, 77)	0.44	0.16 = 0.11 + 0.01 + 0.03	0.41 (32, 89)	0.43
	PA	0.56	6.06	1.25	0.38 = 0.33 + 0.06	1.42 (11, 56)	1.18	0.46 = 0.31 + 0.06 + 0.08	0.28 (25, 99)	1.15
	UK	OOM	OOM	OOM	0.19 = 0.19 + 0.00	4.49 (0, 0)	1.08	0.26 = 0.18 + 0.02 + 0.06	1.39 (18, 72)	1.01
PSG	PR	0.40	1.64	1.75	0.16 = 0.16 + 0.01	0.03 (100, 100)	1.74	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	1.72
	TW	0.72	5.22	2.59	0.23 = 0.22 + 0.02	1.12 (4, 60)	2.60	0.28 = 0.21 + 0.02 + 0.05	0.51 (26, 86)	2.52
	PA	1.86	4.85	5.78	0.54 = 0.49 + 0.05	1.68 (6, 37)	6.09	0.61 = 0.47 + 0.04 + 0.09	0.33 (22, 97)	6.01
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.65 = 0.49 + 0.03 + 0.13	3.37 (13, 57)	7.00

Performance Breakdown



45

S, **E**, and **T** represent Sample, Extract, and Train stages. **G**, **M**, and **C** represent graph sampling, marking cached vertices, and copying samples to host memory in Sample stage, respectively. **R%** and **H%** represent the cache ratio of features and the cache hit rate.

GNN	Dataset	DGL			T_{SOTA}			GNNLab		
		S	E	T	S = G + M	E (R%, H%)	T	S = G + M + C	E (R%, H%)	T
GCN	PR	0.35	2.81	1.22	0.30 = 0.29 + 0.01	0.04 (100, 100)	1.18	0.39 = 0.29 + 0.01 + 0.09	0.15 (100, 100)	1.18
	TW	0.74	9.44	1.48	0.29 = 0.26 + 0.03	3.68 (1, 29)	1.53	0.37 = 0.26 + 0.03 + 0.08	0.76 (25, 89)	1.51
	PA	1.20	10.70	4.00	0.79 = 0.70 + 0.10	3.64 (7, 38)	4.00	0.96 = 0.68 + 0.10 + 0.18	0.49 (21, 99)	3.82
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.56 = 0.39 + 0.03 + 0.14	3.06 (14, 70)	3.09
GSG	PR	0.13	1.92	0.23	0.16 = 0.15 + 0.01	0.03 (100, 100)	0.25	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	0.24
	TW	0.38	4.65	0.44	0.12 = 0.11 + 0.01	0.62 (15, 77)	0.44	0.16 = 0.11 + 0.01 + 0.03	0.41 (32, 89)	0.43
	PA	0.56	6.06	1.25	0.38 = 0.33 + 0.06	1.42 (11, 56)	1.18	0.46 = 0.31 + 0.06 + 0.08	0.28 (25, 99)	1.15
	UK	OOM	OOM	OOM	0.19 = 0.19 + 0.00	4.49 (0, 0)	1.08	0.26 = 0.18 + 0.02 + 0.06	1.39 (18, 72)	1.01
PSG	PR	0.40	1.64	1.75	0.16 = 0.16 + 0.01	0.03 (100, 100)	1.74	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	1.72
	TW	0.72	5.22	2.59	0.23 = 0.22 + 0.02	1.12 (4, 60)	2.60	0.28 = 0.21 + 0.02 + 0.05	0.51 (26, 86)	2.52
	PA	1.86	4.85	5.78	0.54 = 0.49 + 0.05	1.68 (6, 37)	6.09	0.61 = 0.47 + 0.04 + 0.09	0.33 (22, 97)	6.01
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.65 = 0.49 + 0.03 + 0.13	3.37 (13, 57)	7.00

Performance Breakdown



S, **E**, and **T** represent Sample, Extract, and Train stages. **G**, **M**, and **C** represent graph sampling, marking cached vertices, and copying samples to host memory in Sample stage, respectively. **R%** and **H%** represent the cache ratio of features and the cache hit rate.

GNN	Dataset	DGL			T _{SOTA}			GNNLab		
		<u>S</u>	<u>E</u>	<u>T</u>	<u>S</u> = G + M	<u>E</u> (R%, H%)	<u>T</u>	<u>S</u> = G + M + C	<u>E</u> (R%, H%)	<u>T</u>
GCN	PR	0.35	2.81	1.22	0.30 = 0.29 + 0.01	0.04 (100, 100)	1.18	0.39 = 0.29 + 0.01 + 0.09	0.15 (100, 100)	1.18
	TW	0.74	9.44	1.48	0.29 = 0.26 + 0.03	3.68 (1, 29)	1.53	0.37 = 0.26 + 0.03 + 0.08	0.76 (25, 89)	1.51
	PA	1.20	10.70	4.00	0.79 = 0.70 + 0.10	3.64 (7, 38)	4.00	0.96 = 0.68 + 0.10 + 0.18	0.49 (21, 99)	3.82
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.56 = 0.39 + 0.03 + 0.14	3.06 (14, 70)	3.09
GSG	PR	0.13	1.92	0.23	0.16 = 0.15 + 0.01	0.03 (100, 100)	0.25	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	0.24
	TW	0.38	4.65	0.44	0.12 = 0.11 + 0.01	0.62 (15, 77)	0.44	0.16 = 0.11 + 0.01 + 0.03	0.41 (32, 89)	0.43
	PA	0.56	6.06	1.25	0.38 = 0.33 + 0.06	1.42 (11, 56)	1.18	0.46 = 0.31 + 0.06 + 0.08	0.28 (25, 99)	1.15
	UK	OOM	OOM	OOM	0.19 = 0.19 + 0.00	4.49 (0, 0)	1.08	0.26 = 0.18 + 0.02 + 0.06	1.39 (18, 72)	1.01
PSG	PR	0.40	1.64	1.75	0.16 = 0.16 + 0.01	0.03 (100, 100)	1.74	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	1.72
	TW	0.72	5.22	2.59	0.23 = 0.22 + 0.02	1.12 (4, 60)	2.60	0.28 = 0.21 + 0.02 + 0.05	0.51 (26, 86)	2.52
	PA	1.86	4.85	5.78	0.54 = 0.49 + 0.05	1.68 (6, 37)	6.09	0.61 = 0.47 + 0.04 + 0.09	0.33 (22, 97)	6.01
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.65 = 0.49 + 0.03 + 0.13	3.37 (13, 57)	7.00

Performance Breakdown



S, **E**, and **T** represent Sample, Extract, and Train stages. **G**, **M**, and **C** represent graph sampling, marking cached vertices, and copying samples to host memory in Sample stage, respectively. **R%** and **H%** represent the cache ratio of features and the cache hit rate.

GNN	Dataset	DGL			T _{SOTA}			GNNLab		
		<u>S</u>	<u>E</u>	<u>T</u>	<u>S</u> = G + M	<u>E</u> (R%, H%)	<u>T</u>	<u>S</u> = G + M + C	<u>E</u> (R%, H%)	<u>T</u>
GCN	PR	0.35	2.81	1.22	0.30 = 0.29 + 0.01	0.04 (100, 100)	1.18	0.39 = 0.29 + 0.01 + 0.09	0.15 (100, 100)	1.18
	TW	0.74	9.44	1.48	0.29 = 0.26 + 0.03	3.68 (1, 29)	1.53	0.37 = 0.26 + 0.03 + 0.08	0.76 (25, 89)	1.51
	PA	1.20	10.70	4.00	0.79 = 0.70 + 0.10	3.64 (7, 38)	4.00	0.96 = 0.68 + 0.10 + 0.18	0.49 (21, 99)	3.82
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.56 = 0.39 + 0.03 + 0.14	3.06 (14, 70)	3.09
GSG	PR	0.13	1.92	0.23	0.16 = 0.15 + 0.01	0.03 (100, 100)	0.25	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	0.24
	TW	0.38	4.65	0.44	0.12 = 0.11 + 0.01	0.62 (15, 77)	0.44	0.16 = 0.11 + 0.01 + 0.03	0.41 (32, 89)	0.43
	PA	0.56	6.06	1.25	0.38 = 0.33 + 0.06	1.42 (11, 56)	1.18	0.46 = 0.31 + 0.06 + 0.08	0.28 (25, 99)	1.15
	UK	OOM	OOM	OOM	0.19 = 0.19 + 0.00	4.49 (0, 0)	1.08	0.26 = 0.18 + 0.02 + 0.06	1.39 (18, 72)	1.01
PSG	PR	0.40	1.64	1.75	0.16 = 0.16 + 0.01	0.03 (100, 100)	1.74	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	1.72
	TW	0.72	5.22	2.59	0.23 = 0.22 + 0.02	1.12 (4, 60)	2.60	0.28 = 0.21 + 0.02 + 0.05	0.51 (26, 86)	2.52
	PA	1.86	4.85	5.78	0.54 = 0.49 + 0.05	1.68 (6, 37)	6.09	0.61 = 0.47 + 0.04 + 0.09	0.33 (22, 97)	6.01
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.65 = 0.49 + 0.03 + 0.13	3.37 (13, 57)	7.00

Performance Breakdown



S, **E**, and **T** represent Sample, Extract, and Train stages. **G**, **M**, and **C** represent graph sampling, marking cached vertices, and copying samples to host memory in Sample stage, respectively. **R%** and **H%** represent the cache ratio of features and the cache hit rate.

GNN	Dataset	DGL			T _{SOTA}			GNNLab		
		S	E	T	S = G + M	E (R%, H%)	T	S = G + M + C	E (R%, H%)	T
GCN	PR	0.35	2.81	1.22	0.30 = 0.29 + 0.01	0.04 (100, 100)	1.18	0.39 = 0.29 + 0.01 + 0.09	0.15 (100, 100)	1.18
	TW	0.74	9.44	1.48	0.29 = 0.26 + 0.03	3.68 (1, 29)	1.53	0.37 = 0.26 + 0.03 + 0.08	0.76 (25, 89)	1.51
	PA	1.20	10.70	4.00	0.79 = 0.70 + 0.10	3.64 (7, 38)	4.00	0.96 = 0.68 + 0.10 + 0.18	0.49 (21, 99)	3.82
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.56 = 0.39 + 0.03 + 0.14	3.06 (14, 70)	3.09
GSG	PR	0.13	1.92	0.23	0.16 = 0.15 + 0.01	0.03 (100, 100)	0.25	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	0.24
	TW	0.38	4.65	0.44	0.12 = 0.11 + 0.01	0.62 (15, 77)	0.44	0.16 = 0.11 + 0.01 + 0.03	0.41 (32, 89)	0.43
	PA	0.56	6.06	1.25	0.38 = 0.33 + 0.06	1.42 (11, 56)	1.18	0.46 = 0.31 + 0.06 + 0.08	0.28 (25, 99)	1.15
	UK	OOM	OOM	OOM	0.19 = 0.19 + 0.00	4.49 (0, 0)	1.08	0.26 = 0.18 + 0.02 + 0.06	1.39 (18, 72)	1.01
PSG	PR	0.40	1.64	1.75	0.16 = 0.16 + 0.01	0.03 (100, 100)	1.74	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	1.72
	TW	0.72	5.22	2.59	0.23 = 0.22 + 0.02	1.12 (4, 60)	2.60	0.28 = 0.21 + 0.02 + 0.05	0.51 (26, 86)	2.52
	PA	1.86	4.85	5.78	0.54 = 0.49 + 0.05	1.68 (6, 37)	6.09	0.61 = 0.47 + 0.04 + 0.09	0.33 (22, 97)	6.01
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.65 = 0.49 + 0.03 + 0.13	3.37 (13, 57)	7.00

Performance Breakdown



49

S, **E**, and **T** represent Sample, Extract, and Train stages. **G**, **M**, and **C** represent graph sampling, marking cached vertices, and copying samples to host memory in Sample stage, respectively. **R%** and **H%** represent the cache ratio of features and the cache hit rate.

GNN	Dataset	DGL			T _{SOTA}			GNNLab		
		<u>S</u>	<u>E</u>	<u>T</u>	<u>S</u> = G + M	<u>E</u> (R%, H%)	<u>T</u>	<u>S</u> = G + M + C	<u>E</u> (R%, H%)	<u>T</u>
GCN	PR	0.35	2.81	1.22	0.30 = 0.29 + 0.01	0.04 (100, 100)	1.18	0.39 = 0.29 + 0.01 + 0.09	0.15 (100, 100)	1.18
	TW	0.74	9.44	1.48	0.29 = 0.26 + 0.03	3.68 (1, 29)	1.53	0.37 = 0.26 + 0.03 + 0.08	0.76 (25, 89)	1.51
	PA	1.20	10.70	4.00	0.79 = 0.70 + 0.10	3.64 (7, 38)	4.00	0.96 = 0.68 + 0.10 + 0.18	0.49 (21, 99)	3.82
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.56 = 0.39 + 0.03 + 0.14	3.06 (14, 70)	3.09
GSG	PR	0.13	1.92	0.23	0.16 = 0.15 + 0.01	0.03 (100, 100)	0.25	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	0.24
	TW	0.38	4.65	0.44	0.12 = 0.11 + 0.01	0.62 (15, 77)	0.44	0.16 = 0.11 + 0.01 + 0.03	0.41 (32, 89)	0.43
	PA	0.56	6.06	1.25	0.38 = 0.33 + 0.06	1.42 (11, 56)	1.18	0.46 = 0.31 + 0.06 + 0.08	0.28 (25, 99)	1.15
	UK	OOM	OOM	OOM	0.19 = 0.19 + 0.00	4.49 (0, 0)	1.08	0.26 = 0.18 + 0.02 + 0.06	1.39 (18, 72)	1.01
PSG	PR	0.40	1.64	1.75	0.16 = 0.16 + 0.01	0.03 (100, 100)	1.74	0.20 = 0.15 + 0.01 + 0.04	0.08 (100, 100)	1.72
	TW	0.72	5.22	2.59	0.23 = 0.22 + 0.02	1.12 (4, 60)	2.60	0.28 = 0.21 + 0.02 + 0.05	0.51 (26, 86)	2.52
	PA	1.86	4.85	5.78	0.54 = 0.49 + 0.05	1.68 (6, 37)	6.09	0.61 = 0.47 + 0.04 + 0.09	0.33 (22, 97)	6.01
	UK	OOM	OOM	OOM	OOM	OOM	OOM	0.65 = 0.49 + 0.03 + 0.13	3.37 (13, 57)	7.00

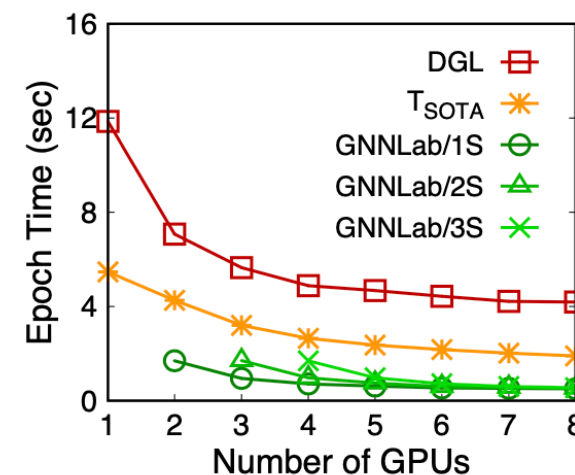
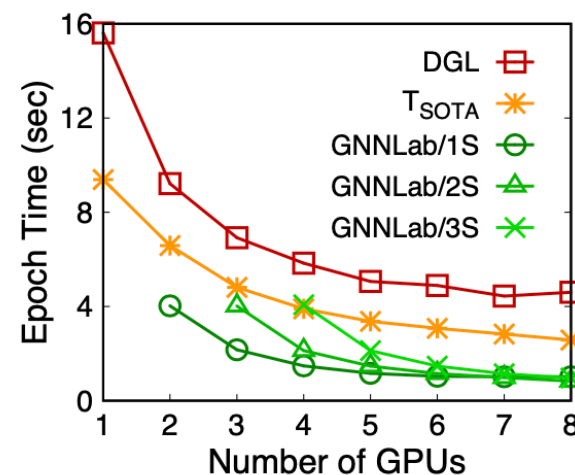
Scalability



50

- ▶ DGL and T_{SOTA}
 - ▶ Time sharing design
 - ▶ More work on CPUs
 - ▶ CPUs stop growing

GCN on PA and TW



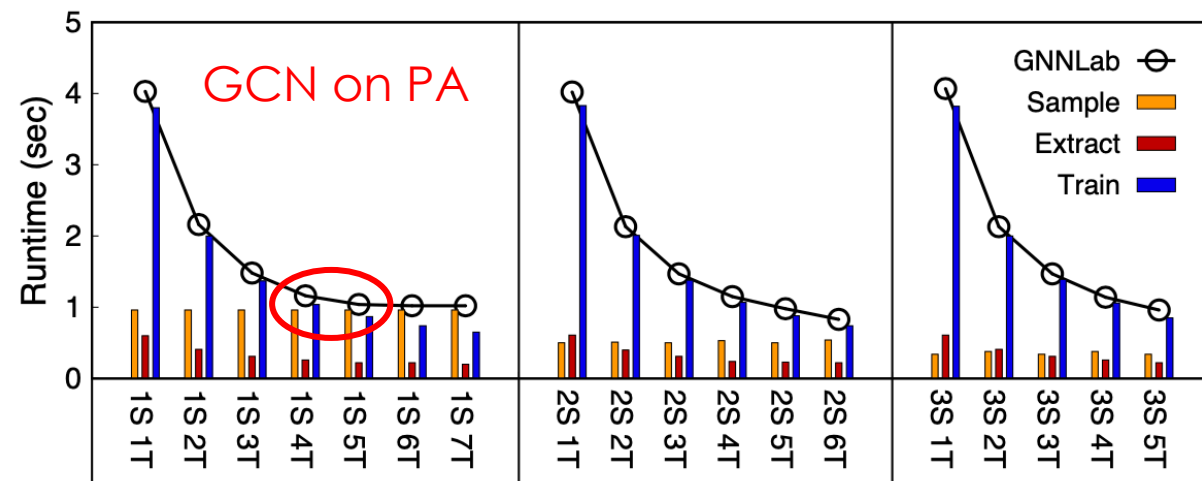
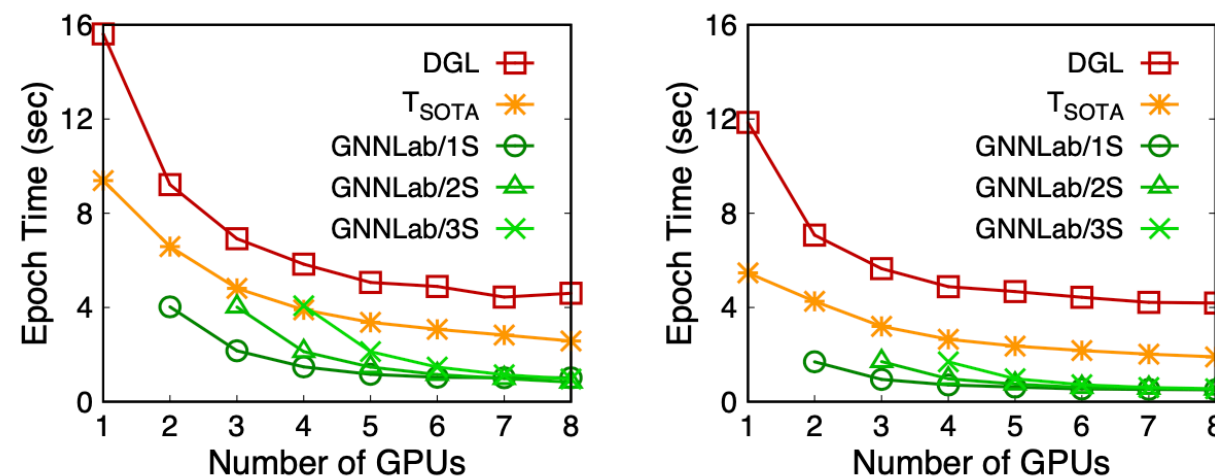
Scalability



51

- ▶ DGL and T_{SOTA}
 - ▶ Time sharing design
 - ▶ More work on CPUs
 - ▶ CPUs stop growing
- ▶ GNNLab
 - ▶ Good parallelism
 - ▶ Bottleneck may change
 - ▶ 1 Sampler is **not enough**
 - ▶ 3 Sampler is **too much**

GCN on PA and TW



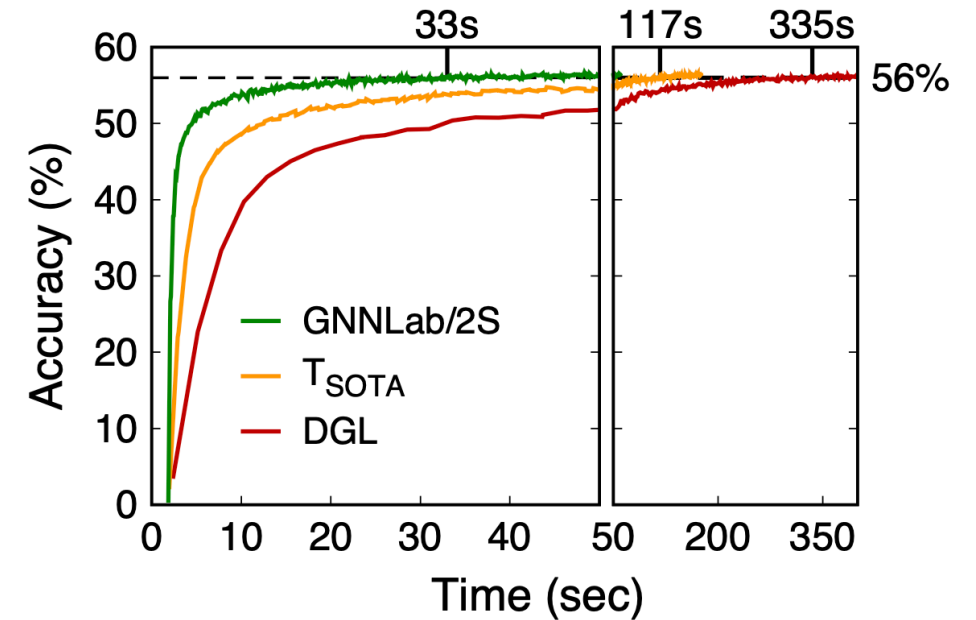
Training Convergence



52

Case: GraphSAGE on Papers100M

- ▶ Converge to same accuracy targets
- ▶ GNNLab outperforms
DGL by $10.2\times$ and T_{SOTA} by $3.5\times$



Training Convergence



53

Case: GraphSAGE on Papers100M

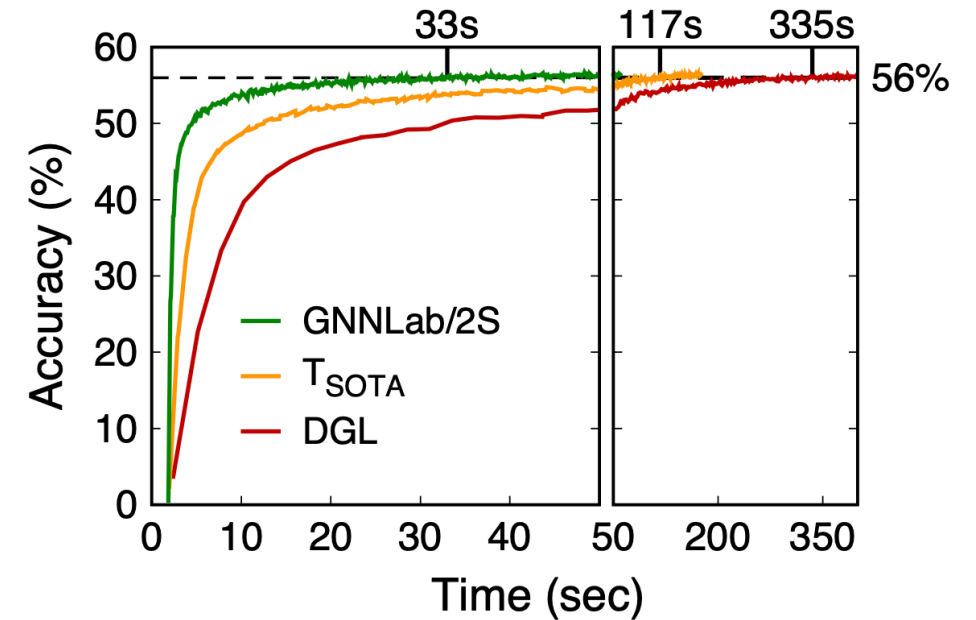
► Converge to same accuracy targets

► GNNLab outperforms

DGL by $10.2\times$ and T_{SOTA} by $3.5\times$

1. **Faster** training (per epoch)

► vs. DGL by $8.2\times$ and T_{SOTA} by $2.8\times$



Training Convergence



54

Case: GraphSAGE on Papers100M

- ▶ Converge to same accuracy targets
- ▶ GNNLab outperforms

DGL by $10.2\times$ and T_{SOTA} by $3.5\times$

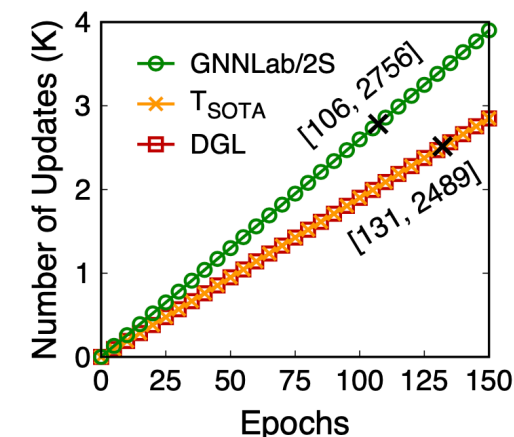
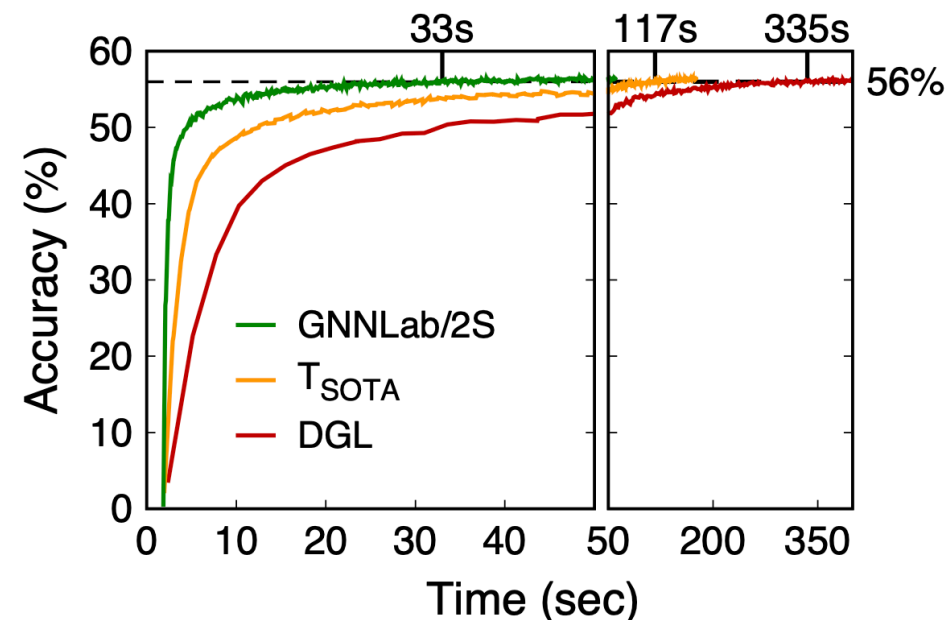
1. **Faster** training (per epoch)

- ▶ vs. DGL by $8.2\times$ and T_{SOTA} by $2.8\times$

2. **Fewer** epochs, reduced by $1.24\times$

- ▶ GNNLab: **106** (6 GPU workers for training)
- ▶ DGL/ T_{SOTA} : **131** (8 GPU workers for training)

The **more** GPUs allocated for model training, the **fewer** gradient updates per training epoch, and **more** epochs are required to achieve the same expected accuracy.



Conclusion & Thanks



55

GNNLab: a factored system for sample-based GNN training

- ▶ Replace time sharing with **space sharing** design
- ▶ **Flexible** architecture and scheduling for load balance
- ▶ A new **efficient** and **robust** caching policy



GNNLab will be published in **EuroSys 2022**

Artifact Evaluation: <https://github.com/SJTU-IPADS/fgnn-artifacts>

Open source: <https://github.com/SJTU-IPADS/gnnlab> (available soon)

AI needs Systems Research

Space sharing advance

- ▶ Fine-grained: inside a GPU (MPS and MIG)
- ▶ Decouple GPU CU and GPU memory

How about other GNNs and sampling algorithms?

- ▶ ClusterGCN and Shallow Subgraph Samplers

How to contribute to DGL?