

VertexFlow: Dynamic Scheduling for Heterogeneous Edge Inference

Lightning Talk & Poster Teaser (Poster #304)

Authors: Sarah Jenkins David K. Sterling

Affiliations: Synapse Institute

Apex Labs

 letx.app/m/vertexflow

 {jenkins, sterling}@synapse.org

Key Result: Heterogeneous Task Stealing

Outperforming Homogeneous Pipelines

Dynamic Workload Balancing

- Edge nodes exhibit transient resource fluctuations.
- Homogeneous pipelines suffer from **head-of-line blocking**.
- VertexFlow introduces a zero-overhead task stealing protocol across CPU/GPU cores.

Performance Gain: We achieve a $5.4\times$ speedup in end-to-end inference throughput under dynamic load [1].

Throughput Comparison (FPS)

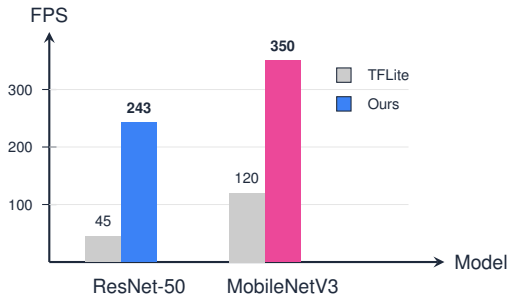


Figure: VertexFlow throughput vs. standard TFLite runtime.

Join Us at Poster #304

Let's Discuss the Details

Poster Session II

- Time: Tuesday, 14:00 – 17:30
- Location: Hall B, Section C
- Poster Number: **#304**

Discussion Highlights:

1. Mathematical proofs for task stealing overhead bounds.
2. Portability across ARM Cortex-M and RISC-V targets.
3. Demonstration on actual hardware prototypes.

Scan for Paper & Code



 github.com/synapse-labs
 letx.app/m/vertexflow