

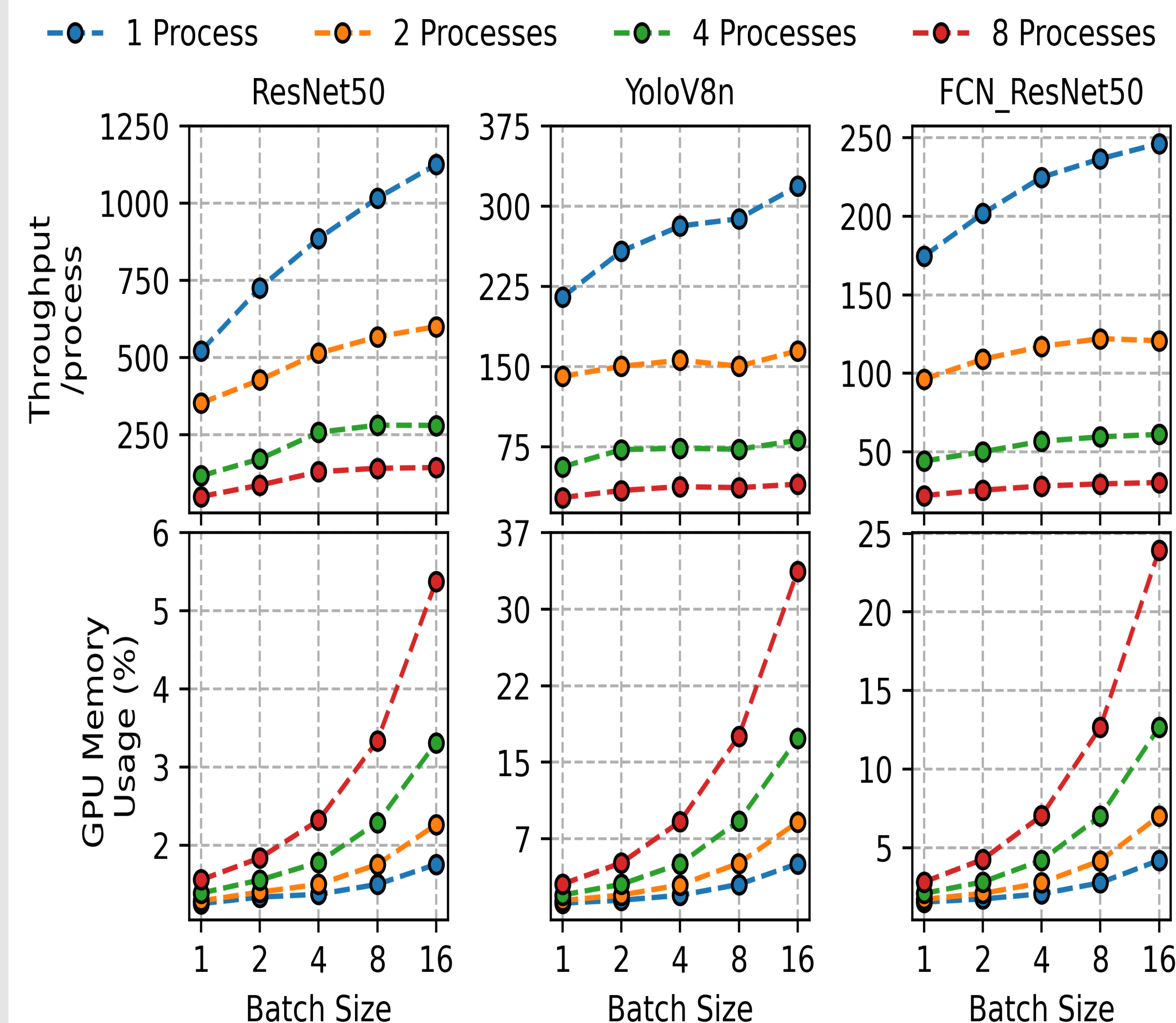
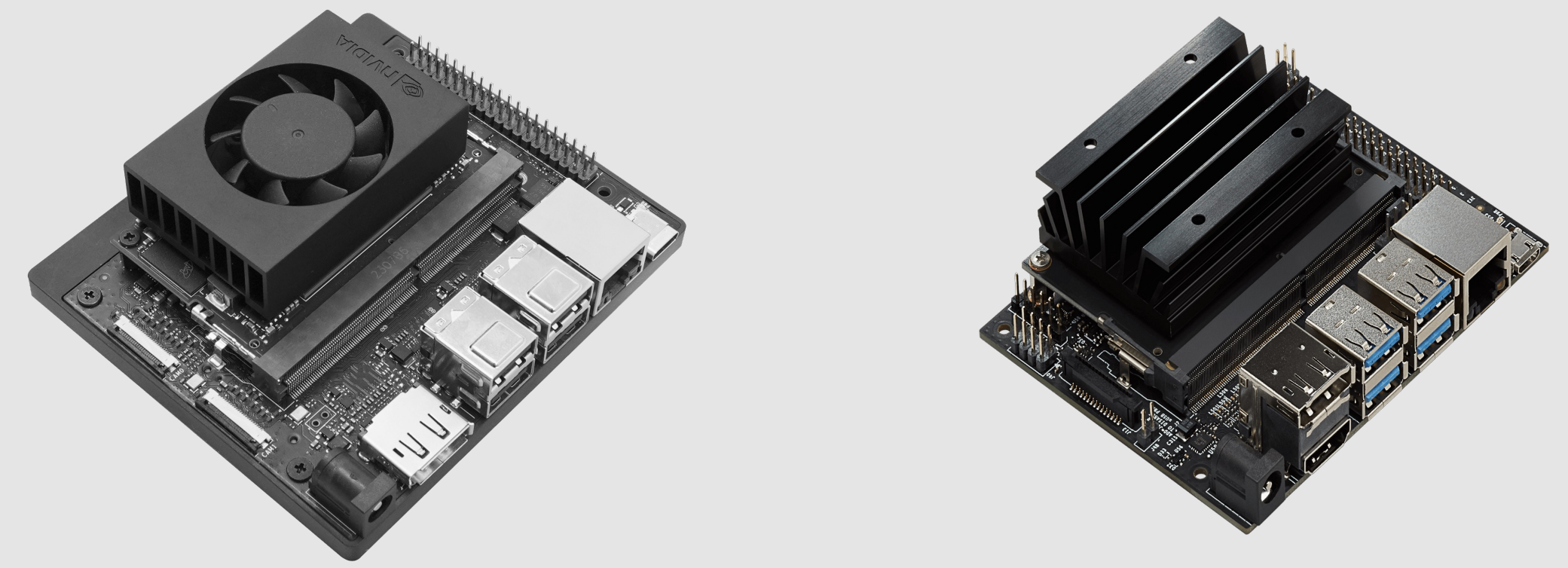
NETMODEL GROUP - INTERNET AND DATA SCIENCE LAB

Abhinaba Chakraborty, Didier Colle, Mario Pickavet, Akis Kourtis, Andreas Oikonomakis, Wouter Tavernier

PROFILING CONCURRENT VISION INFERENCE WORKLOADS ON NVIDIA JETSON

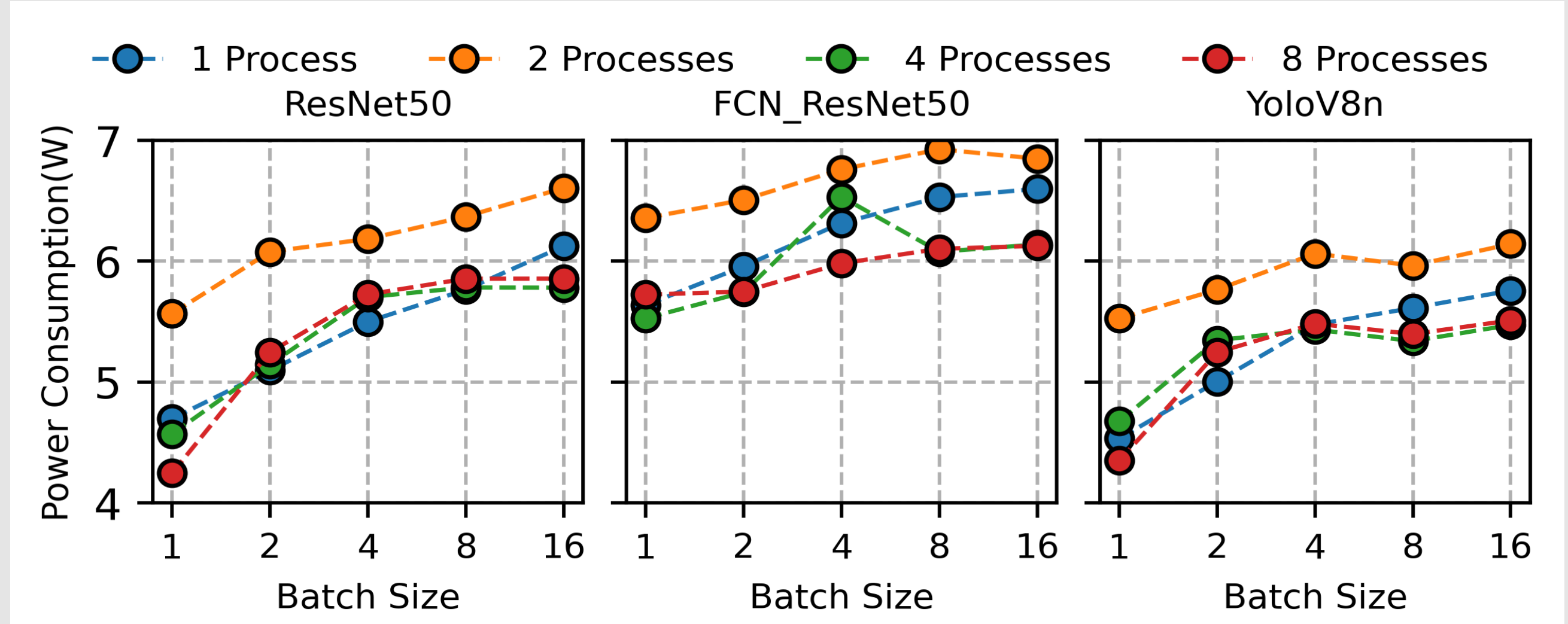
System Throughput vs Number of Process

$process \uparrow \text{throughput} / process \downarrow GPU \text{ Memory Usage} \uparrow$
 $batch \uparrow \text{throughput} / process \uparrow GPU \text{ Memory Usage} \uparrow$



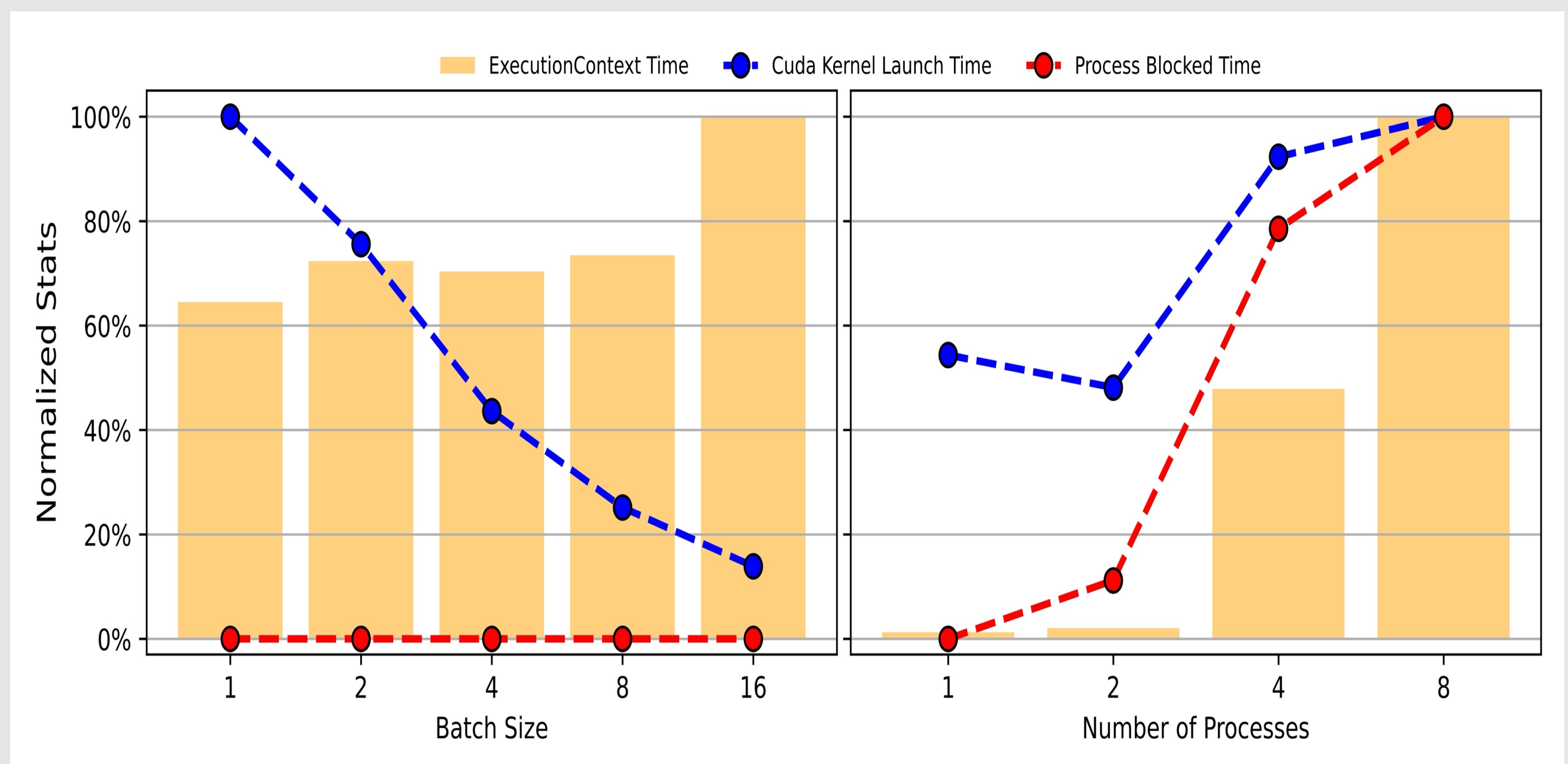
Power Consumption vs Number of Process

$power \text{ consumption} \propto \text{throughput}$



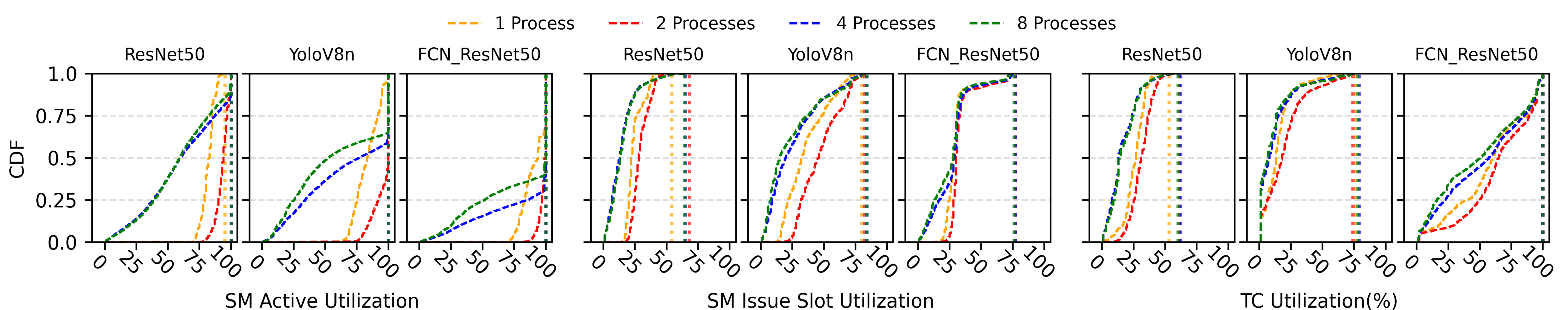
Low Level Performance Metrics

$process \uparrow EC \text{ Time} \uparrow \text{Kernel Launch Time} \uparrow$
 $Batch \uparrow \text{Kernel Launch Time} \downarrow$



Kernel Level Performance Analysis

$SM \text{ Utilization} \sim 80 - 100\%$
 $Issue \text{ Slot Utilization} \sim 25\%$
 $Tensor \text{ Core Utilization} \sim 25\%$



Conclusions

1. Need for Fine-Grained CPU-GPU Co-Scheduling of Threads
2. Development of orchestration techniques for concurrent DL executions across CPU-GPU

Contact

abhinaba.chakraborty@ugent.be
<https://www.idlab.ugent.be/research-teams/netmodel>

Universiteit Gent

@ugent

Ghent University