

Auto-reconfiguration for Latency Minimization in CPU-based DNN Serving

Ankit Bhardwaj
MIT CSAIL

Amar Phanishayee
Meta

Deepak Narayanan
NVIDIA

Ryan Stutsman
University of Utah

Motivation

Serving small DNN models on CPU:

- Consumes lower power
- Cost-effective
- Easy to scale across multiple machines

Increasing core counts and memory bandwidth make them suitable for serving small models.

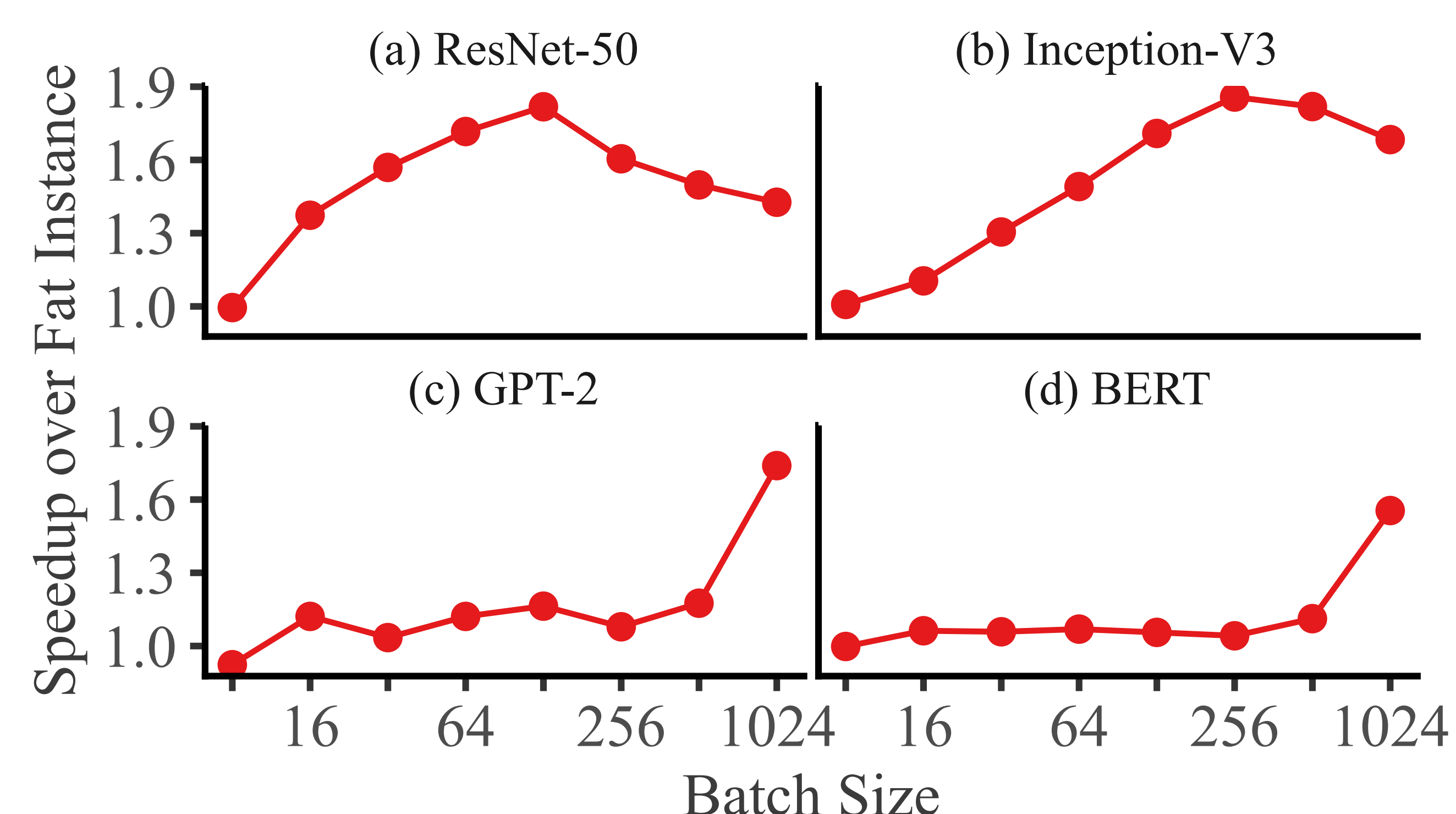
Identifying Configuration is Challenging

$\langle 1, T, B \rangle$ to $\langle i, t, b \rangle$ is challenging

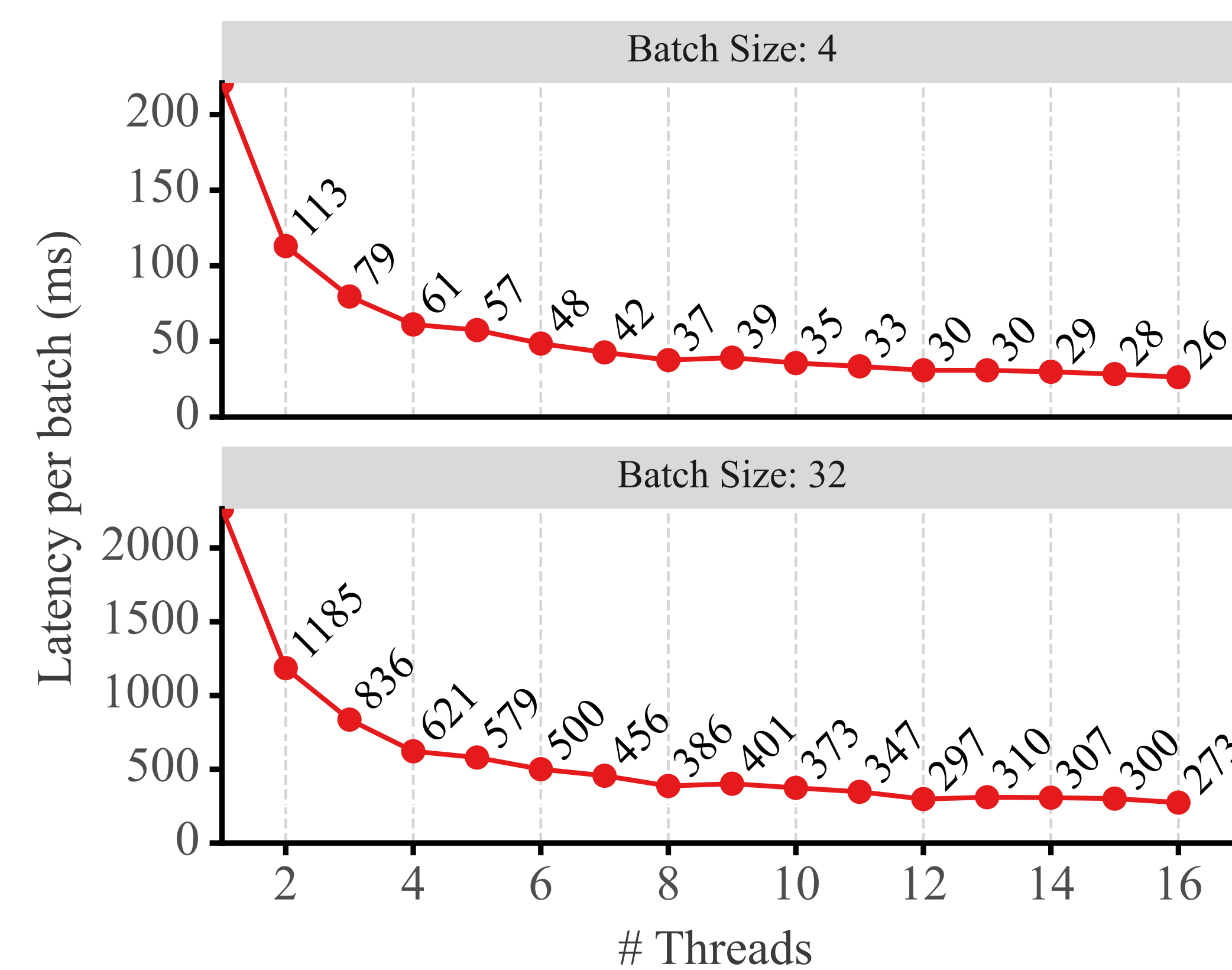
- Model specific parallelism
- Batch-specific gains
- Heterogeneous hardware

Identifying a configuration that works for a specific model and a given batch size on a machine can be challenging.

Multi-instance Speedups

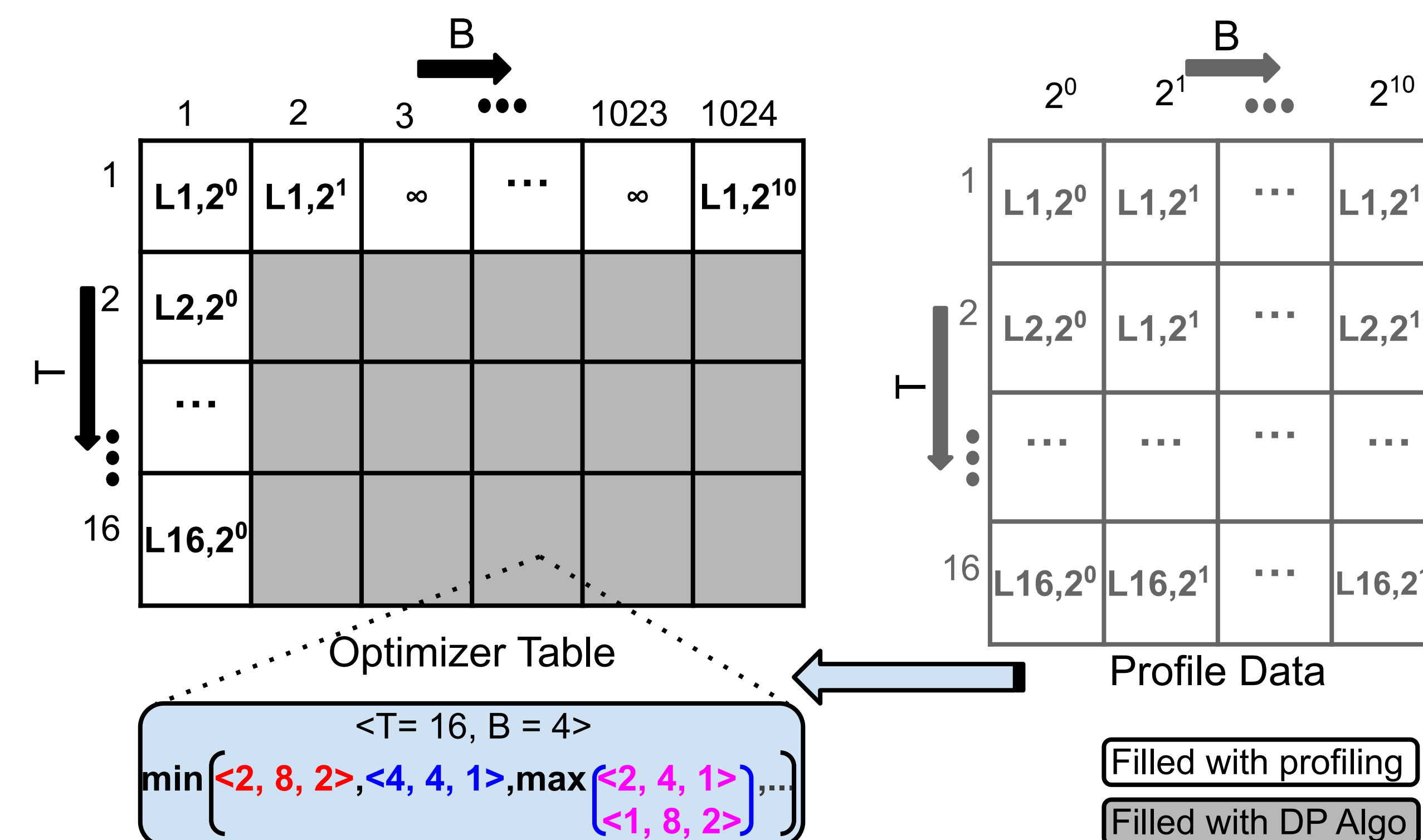


Reducing Latency with Intra-op Parallelism



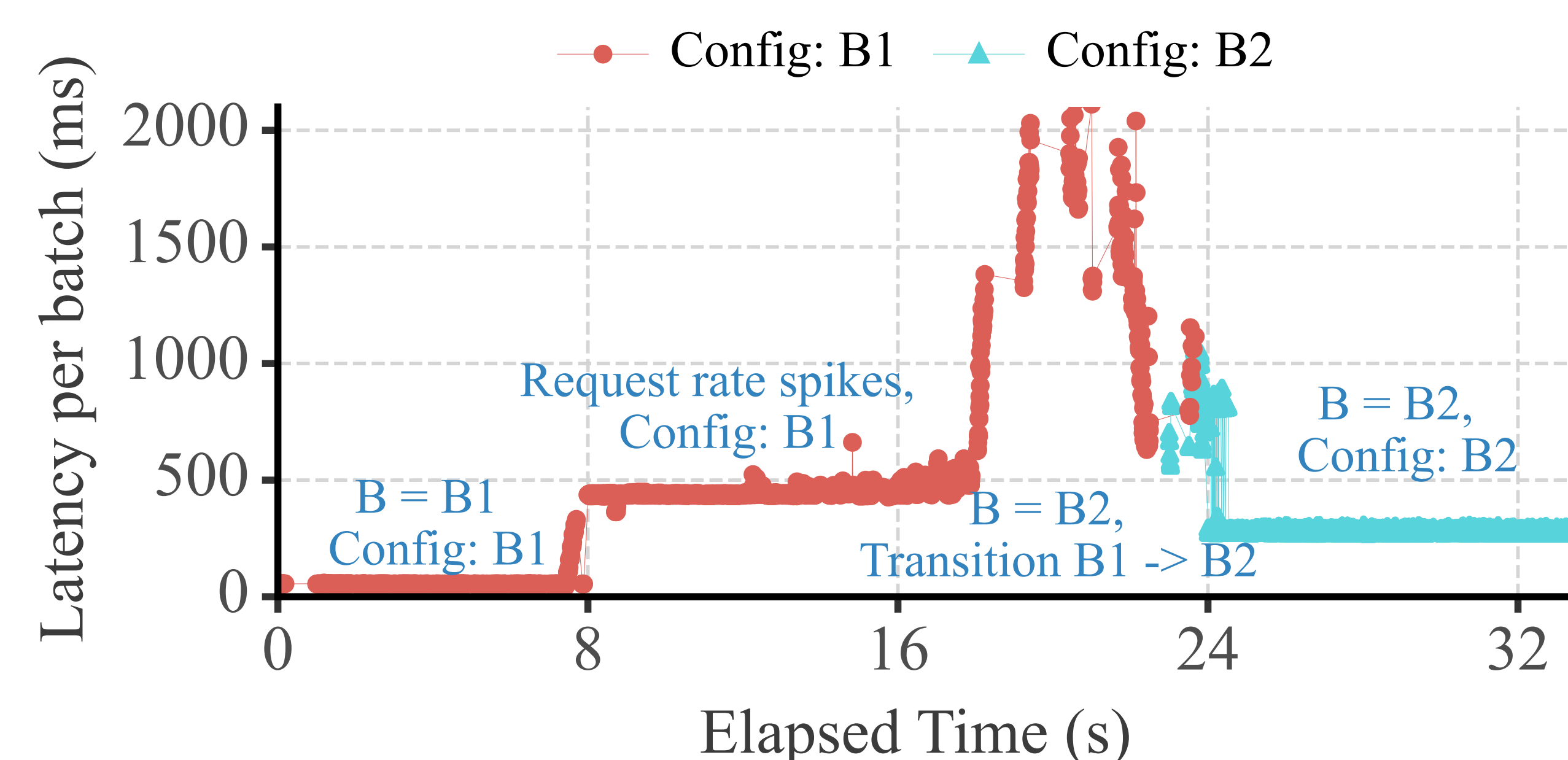
Diminishing returns on adding more cores

Dynamic Programming to the Rescue

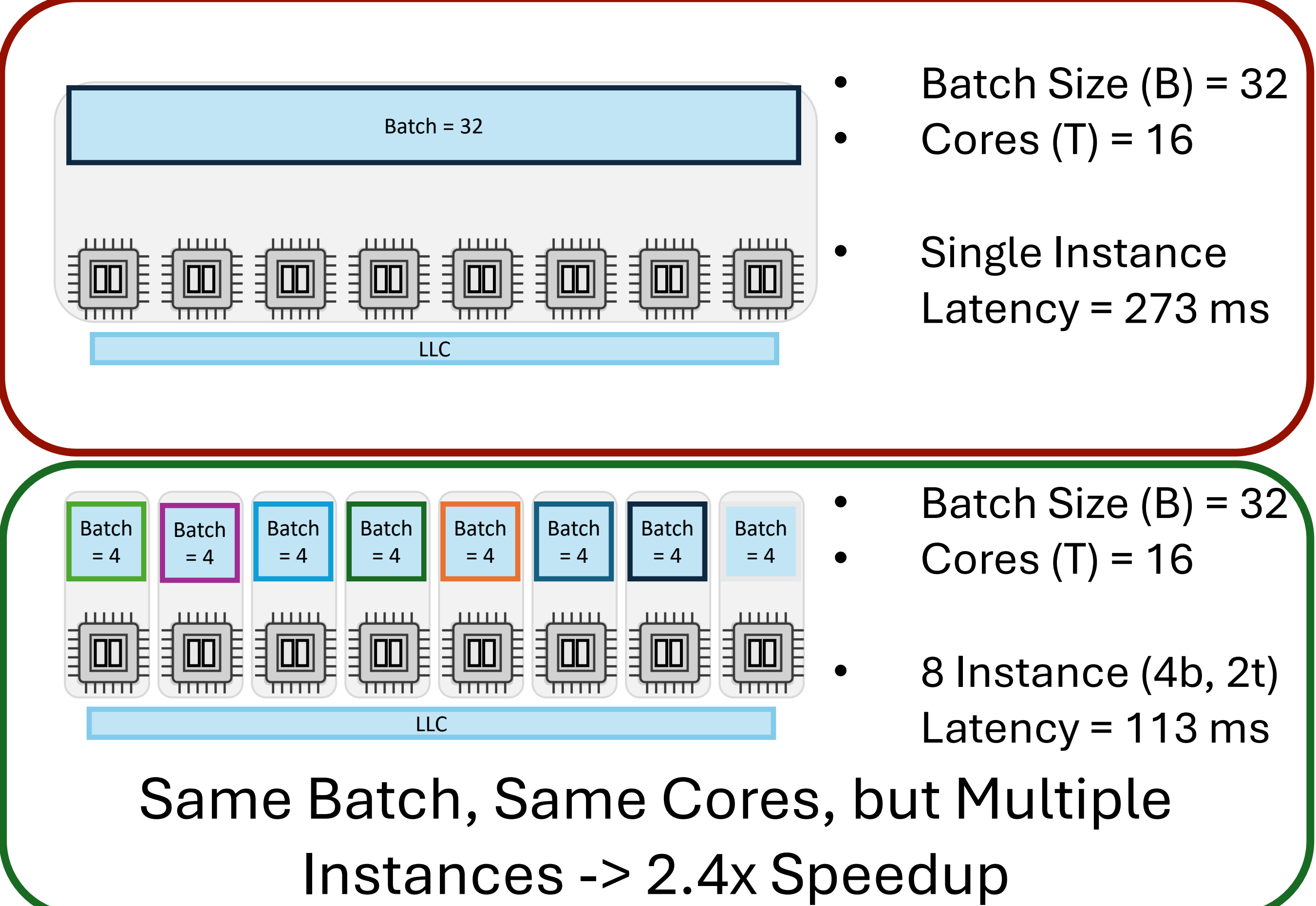


- Profile a few configurations for each model
- Predict the rest using a 2D-Knapsack solution

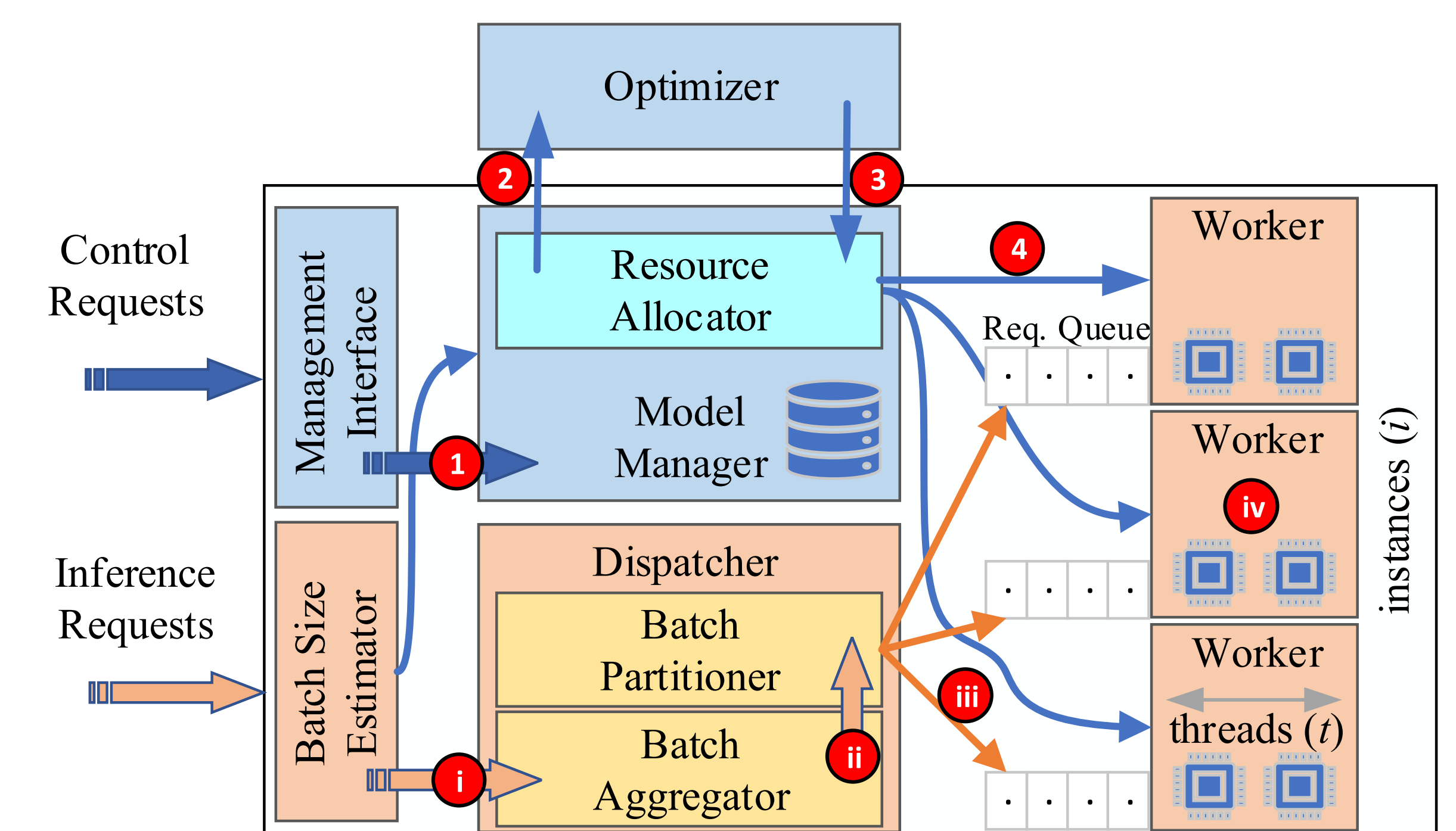
Automated Reconfiguration



Multi-instance Serving for Lower Latency



Packrat Architecture



Conclusion

- Pure intra-op parallelism on CPUs shows diminishing returns.
- Multi-instance execution improves performance.
- Packrat automatically finds optimal configurations and dynamically adjusts the serving system.
- Achieves up to **3.2x** latency speedup.