

Latency-Aware Task Scheduling in Heterogeneous Edge-Cloud Continuum: An Empirical Analysis of Preemptive DAG Partitioning Under Stochastic Workload Volatility

Jesse Anderson

Professor

Korea Advanced Institute of Science and Technology (KAIST)
291 Daehak-ro, Yuseong-gu, Daejeon 34141, Republic of Korea

Casey Campbell

Associate Professor

Technische Universität München
Arcisstraße 21, 80333 Munich, Bavaria, Germany

Cameron Harris

PhD

University of Waterloo
200 University Avenue West, Waterloo, Ontario N2L 3G1, Canada

Robin Carter

D.Sc

Delft University of Technology
Mekelweg 2, 2628 CD Delft, South Holland, Netherlands

Abstract. Edge-cloud continuum architectures present non-trivial scheduling challenges when directed acyclic graph (DAG)-structured computational workloads exhibit stochastic inter-task dependency volatility. This paper presents an empirical analysis of preemptive DAG partitioning strategies deployed across heterogeneous edge nodes and cloud back-ends, evaluating end-to-end latency, resource utilization efficiency, and fault-tolerance resilience. We instrument a real-world testbed comprising twelve heterogeneous edge devices and a multi-tenant cloud cluster, subjecting it to synthetically generated and trace-driven workloads sampled from production microservice pipelines. Experimental results demonstrate that a latency-aware, criticality-weighted partitioning heuristic reduces makespan by 23.7% and tail latency (P99) by 31.4% over baseline First-Fit and HEFT schedulers, while maintaining energy overhead within acceptable operational margins. These findings establish quantifiable benchmarks for edge-cloud co-scheduling in latency-sensitive industrial IoT deployments.

Keywords: edge-cloud continuum, DAG task scheduling, heterogeneous computing, preemptive partitioning, stochastic workload modeling, makespan optimization, tail latency, industrial IoT orchestration, criticality-weighted heuristics

Introduction

The proliferation of latency-sensitive applications—spanning autonomous vehicular systems, real-time industrial process control, augmented reality pipelines, and federated healthcare monitoring—has fundamentally restructured the computational demands placed upon distributed infrastructure. Traditional cloud-centric paradigms, which offload computation to geographically distant data centers, are increasingly incapable of satisfying the sub-millisecond response requirements mandated by these application classes. The emergence of the edge-cloud continuum as an architectural substrate offers a compelling alternative: a fluid computational fabric in which processing tasks are dynamically distributed across a spectrum of heterogeneous resources, ranging from resource-constrained edge nodes co-located with sensing infrastructure to full-scale cloud back-ends provisioned with elastic compute capacity. Central to the operational efficiency of such architectures is the problem of task scheduling—specifically, the partitioning and assignment of complex, dependency-structured workloads modeled as directed acyclic graphs (DAGs) to heterogeneous execution environments. DAG scheduling in homogeneous settings is itself NP-hard, and the introduction of heterogeneous processing capacities, variable network link bandwidths, and non-deterministic task execution times compounds the complexity substantially. Despite a considerable body of theoretical work on heuristic and meta-heuristic schedulers—including the widely cited Heterogeneous Earliest Finish Time (HEFT) algorithm and its derivatives—empirical validation of these methods under realistic, production-grade stochastic workload volatility remains critically sparse as of 2024–2026. Existing literature predominantly evaluates DAG schedulers under synthetic Poisson-arrival models or stationary workload traces that fail to capture the burstiness, temporal correlation, and dependency-graph morphological variability exhibited by modern microservice orchestration pipelines. Furthermore, the preemptive scheduling dimension—wherein in-flight task execution may be interrupted and migrated to a more suitable execution tier upon detection of latency threshold violations—has received disproportionately little empirical scrutiny in the context of heterogeneous edge-cloud deployments. This gap is particularly consequential in industrial IoT settings, where service-level objective (SLO) violations carry direct operational and safety implications. This paper addresses the foregoing deficit through a rigorous, testbed-driven empirical analysis. We design and implement a latency-aware, criticality-weighted preemptive DAG partitioning framework, designated CW-PreDAG, and evaluate its performance against established baseline schedulers across a controlled heterogeneous infrastructure. The remainder of this paper is organized as follows: Section 2 surveys related work in DAG scheduling, edge-cloud offloading, and stochastic workload characterization. Section 3 formalizes the system model and problem statement, including the DAG representation, heterogeneity model, and stochastic arrival processes. Section 4 details the CW-PreDAG algorithm design, including the criticality estimation module and preemption trigger logic. Section 5 describes the experimental testbed, workload generation methodology, and evaluation metrics. Section 6 presents and analyzes quantitative results, followed by a

discussion of practical deployment implications in Section 7. Section 8 concludes the paper and outlines directions for future investigation.

This is a preliminary version. To read the full version of the article, please purchase a subscription.

References

1. Semeniuk, V. V. (2025). OPTIMIZATION OF LOCAL DEVELOPMENT PROCESS USING DOCKER PHP IMAGE THAT COMES WITH A FULL SET OF TOOLS OUT OF THE BOX: DATABASE AND INTERNATIONALIZATION EXTENSIONS. ВЧЕНІ ЗАПИСКИ, 12025226.