

Mixture-of-Schedulers: An Adaptive Scheduling Agent as a Learned Router for Expert Policies

Xinbo Wang*
Shian Jia*
xinbowang@zju.edu.cn
csjsa@zju.edu.cn
Zhejiang University
Hangzhou, Zhejiang, China

Jing Cao
HangZhou City University
Hangzhou, Zhejiang, China
jingcao@hzcw.edu.cn

Ziyang Huang
Zhejiang University
Hangzhou, Zhejiang, China
3220105926@zju.edu.cn

Mingli Song
Zhejiang University
Hangzhou, Zhejiang, China
brooksong@zju.edu.cn

Abstract

Modern operating system schedulers employ a single, static policy, which struggles to deliver optimal performance across the diverse and dynamic workloads of contemporary systems. This “one-policy-fits-all” approach leads to significant compromises in fairness, throughput, and latency, particularly with the rise of heterogeneous hardware and varied application architectures.

This paper proposes a new paradigm: dynamically selecting the optimal policy from a portfolio of specialized schedulers rather than designing a single, monolithic one. We present the Adaptive Scheduling Agent (ASA), a lightweight framework that intelligently matches workloads to the most suitable “expert” scheduling policy at runtime. ASA’s core is a novel, low-overhead offline/online approach. First, an offline process trains a universal, hardware-agnostic machine learning model to recognize abstract workload patterns from system behaviors. Second, at runtime, ASA continually processes the model’s predictions using a time-weighted probability voting algorithm to identify the workload, then makes a scheduling decision by consulting a pre-configured, machine-specific mapping table to switch to the optimal scheduler via Linux’s `sched_ext` framework. This decoupled architecture allows ASA to adapt to new hardware platforms rapidly without expensive retraining of the core recognition model.

Our evaluation, based on a novel benchmark focused on user-experience metrics, demonstrates that ASA consistently outperforms the default Linux scheduler (EEVDF), achieving superior results in 86.4% of test scenarios. Furthermore, ASA’s selections are near-optimal, ranking among the top three schedulers in 78.6% of all scenarios. This validates our approach as a practical path toward more intelligent, adaptive, and responsive operating system schedulers.

1 Introduction

The design of operating system schedulers has always been an exercise in managing fundamental trade-offs, primarily between fairness, throughput, and latency. For decades, scheduler development has advanced by refining heuristics to find a suitable balance for the hardware and software of the era. However, the complexity of modern systems has amplified these trade-offs to a critical point. The concurrent rise of heterogeneous hardware (e.g., P/E-cores, NUMA, and diverse cache hierarchies) and highly diverse software architectures (e.g., microservices, interactive applications, batch processing) means that a single, fixed set of scheduling rules represents a significant compromise. The evolution of the default Linux scheduler from the $O(1)$ scheduler [28], through the Completely Fair Scheduler (CFS) [33], to the most recent EEVDF [29], while increasingly sophisticated, illustrates that even these advanced designs, engineered for general-purpose robustness, often leave significant performance potential untapped in specific, critical scenarios. The core challenge has thus shifted from merely *balancing* these goals to *dynamically adapting* which goal to prioritize from moment to moment.

In recent years, the research community has pursued this goal of dynamic adaptation through several technical routes. One prominent direction has been the development of *scenario-specific static policies* [11, 22, 25], which essentially pre-select a fixed trade-off, such as prioritizing throughput for datacenter workloads. This approach delivers excellent performance in stable, well-defined environments but lacks the agility required for dynamic, general-purpose systems. Another significant line of work has investigated *AI-assisted dynamic policies* [12, 13, 31], which seek to learn the optimal trade-offs online. While this aligns well with the need for adaptability, the associated costs in terms of runtime overhead, training stability, and model interpretability present practical barriers to widespread adoption [8, 39]. This suggests that while the goal of dynamic adaptation is widely

*Both authors contributed equally to this research.

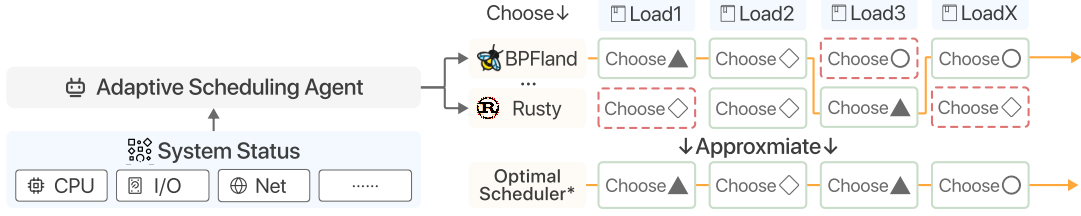


Figure 1. Core Philosophy of ASA: A Paradigm Shift from Designing a Single Optimal Policy to Dynamically Selecting the Best Policy from a Portfolio of Experts.

recognized as correct, the mechanisms for achieving it warrant re-evaluation.

Inspired by the work of *sched_ext* [48], which provides a series of expert schedulers (e.g., LAVD[16] for latency-sensitive workloads), we argue that the key to resolving this dilemma lies in a paradigm shift: from “designing the optimal policy” to “dynamically selecting the optimal policy”. As illustrated in Figure 1, instead of pursuing a single, monolithic scheduler, we propose integrating a portfolio of simpler, specialized schedulers and using an intelligent agent to select the most suitable one at runtime. Based on this philosophy, this research designs and implements the *Adaptive Scheduling Agent* (ASA), a lightweight framework built on the Linux kernel’s extensible scheduler interface, *sched_ext*, which endows the operating system with online awareness and decision-making capabilities.

At its core, ASA operates as an intelligent agent following a classic cycle of *perception*, *decision*, and *action*. In the *perception* module, it continuously monitors runtime metrics to understand the system’s state. The *decision* module then uses a trained model to recognize the active workload pattern and consults a pre-calibrated scheduler mapping table to determine the optimal scheduling strategy. Finally, the *action* module executes this decision by dynamically switching to the most suitable scheduler via the *sched_ext* framework.

The intelligence behind ASA’s decisions is built through a comprehensive three-stage offline preparation process. This process begins with *Prototype Learning* to create a baseline model and scheduler performance map. It is followed by *Dynamic Overhead Calibration*, which simulates scheduler switch and system monitor cost to make ASA take its own operational overhead into account. Finally, *Generalization Model Training* allows a fully operational ASA to refine its models and mappings in a live environment, ensuring high accuracy and adaptability to diverse hardware.

Our work decomposes the complex scheduling decision problem into two more manageable sub-problems: workload pattern recognition and policy matching. This not only avoids the immense challenges of designing a “one-policy-fits-all” scheduler but also, by introducing machine learning, endows the scheduling system with unprecedented flexibility and adaptability.

The main contributions of this research are as follows:

- We design and implement **Adaptive Scheduling Agent (ASA)**, a lightweight, adaptive scheduling agent framework built on eBPF and *sched_ext*, which intelligently routes workloads to expert scheduling policies.
- We propose a **Time-Weighted Probability Voting Algorithm** for workload recognition that uses a voting mechanism with exponential decay to ensure stable and responsive pattern identification, mitigating the impact of transient system noise.
- We introduce a three-stage **Offline Prepare Pipeline** for the ASA that encompasses *Prototype Learning*, *Dynamic Overhead Calibration*, and *Generalization Model Training*, to systematically build a robust scheduling agent.
- We construct a **User Experience Oriented Evaluation Criteria** that measures performance based on metrics directly impacting user experience, such as interaction latency and UI smoothness, bridging the gap left by traditional system-level benchmarks.

Through comprehensive experiments, we have systematically verified the effectiveness of ASA. The experimental results show that ASA achieves a high accuracy in workload pattern recognition and significantly outperforms the default scheduler in the vast majority of test scenarios, approaching the performance of the theoretical optimum. These results validate the effectiveness and practicality of ASA in complex environments, providing a solid and feasible path for building the next generation of intelligent, adaptive operating system kernels.

2 Background & Related Work

This section provides a comprehensive review of the process scheduling landscape, organized into three key areas. We first discuss the foundational challenges that motivate the need for adaptive scheduling, then survey the two dominant scheduler design paradigms: scenario-specific static policies and AI-assisted dynamic policies. Finally, we examine the recent emergence of extensible scheduling frameworks that provide the technical foundation for our work.

2.1 The Enduring Challenge of OS Scheduling

Modern operating system schedulers face a dual challenge, stemming from both technical and methodological issues. First, on the technical front, the prevalence of *heterogeneous hardware architectures*—such as systems with Performance and Efficiency cores (P/E-cores) and Non-Uniform Memory Access (NUMA) designs—complicates resource allocation, as a “one-policy-fits-all” scheduling approach is no longer optimal. Simultaneously, *workload dynamics* have grown in complexity; applications can shift their behavior from I/O-bound to CPU-bound, and interactive processes demand low latency while batch jobs require high throughput. Assigning the right process to the right core at the right time has become a significant challenge for general-purpose schedulers like Linux’s EEVDF [16, 36]. Second, a critical methodological challenge arises from the *limitations of evaluation systems*. Most academic research evaluates schedulers using system-level metrics, such as overall throughput or task completion time [18, 38], which often fail to capture the nuanced performance requirements of different applications. For instance, optimizing for system-wide throughput may degrade the responsiveness of latency-sensitive interactive tasks. Existing benchmarks, like hackbench [5], are often designed to stress test specific kernel mechanisms rather than to simulate the complex, dynamic, and mixed workloads seen in real-world user scenarios, thereby providing an incomplete picture of a scheduler’s true efficacy.

2.2 Scenario-Specific Static Policies

To address these challenges, the more traditional path in scheduler design involves creating *scenario-specific static policies*. This approach treats scheduling as an engineering problem to be solved for a well-defined context, resulting in highly-optimized schedulers for domains like real-time systems [25], virtualized environments [20], and large-scale datacenters [6, 11, 22, 23, 35, 50]. The core idea is to sacrifice generality for performance and low overhead. The design philosophy of specialization persists even as new, flexible kernel frameworks become available, guiding the development of schedulers tailored for specific use cases, such as LAVD for gaming workloads [16]. While effective in their target scenarios, the performance of these static policies can degrade when they encounter dynamic, mixed workloads that fall outside their original design parameters.

2.3 AI-Assisted Dynamic Policies

The second path, driven by advances in machine learning, centers on creating a single, *AI-assisted dynamic policy* capable of adapting to any situation [8, 12, 13, 26, 27, 30, 31, 41]. However, the pursuit of universal adaptability often introduces substantial complexity and overhead. For example,

applying reinforcement learning to task distribution can require significant training phases to be effective [39]. Furthermore, addressing resource management in similarly complex domains, such as mixed database workloads, has led to the exploration of highly sophisticated architectures, including those that combine hierarchical deep reinforcement learning with graph neural networks [49]. This illustrates a trend where increasing adaptability can correspond with a significant rise in model and system complexity.

2.4 The Rise of Extensible Scheduling Frameworks

Extensible scheduling frameworks have emerged as a promising direction to overcome these limitations by decoupling scheduling logic from the kernel, thereby enhancing flexibility and accelerating development. Pioneering work in this area includes Google’s *Ghost* [19] and Stanford’s *Skyloft* [21], which demonstrated the potential of implementing complex scheduling policies outside the core kernel. However, these early frameworks often require significant, non-standard kernel modifications, which complicates their deployment in production environments. A key advancement in this domain is the *sched_ext* framework, which has been accepted into the mainline Linux kernel [48]. As a BPF-based extensible scheduler, *sched_ext* provides a standardized, lightweight, and safer mechanism for deploying and dynamically switching scheduling policies, clearing the path for widespread adoption of flexible scheduling solutions [16].

ASA builds upon the foundation of *sched_ext*, but takes a novel approach. While *sched_ext* provides the *mechanism* for policy switching, it does not provide the *intelligence* for when to switch. Our work fills this critical gap by introducing an adaptive agent that performs workload-pattern recognition and automated policy selection. By leveraging a portfolio of expert schedulers, ASA achieves system-wide adaptability without the high overhead of complex, monolithic AI-driven schedulers, striking a more effective balance between performance and flexibility.

3 ASA: Adaptive Scheduling Agent

3.1 Agent-Based Architecture

The core philosophy of the *Adaptive Scheduling Agent* (ASA) is based on the observation that different workload patterns have fundamentally different optimization requirements. For example, interactive applications prioritize low latency and responsiveness, while batch processing tasks focus on maximizing throughput. Rather than compromising between these conflicting requirements, ASA maintains a portfolio of specialized schedulers and intelligently switches between them as workload patterns change.

3.2 Perception module

The foundation of ASA’s intelligence lies in its ability to accurately perceive and characterize system workloads. The

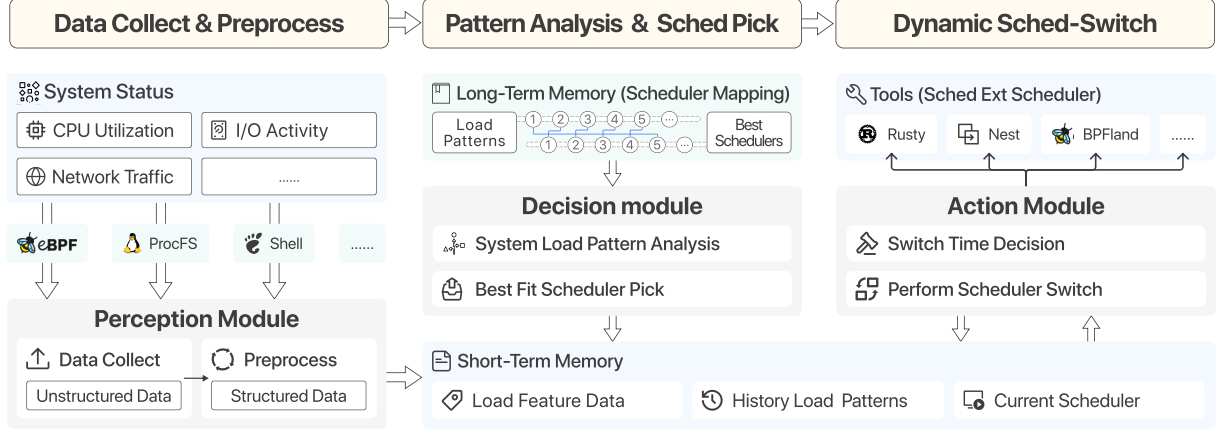


Figure 2. Core Framework of the Adaptive Scheduling Agent

Table 1. the System State Metrics that ASA monitors

Category	Metrics
CPU	User / System Mode Utilization, Nice / Idle / I/O Wait / Hardware Interrupt / Software Interrupt / Steal Utilization, Hot-Spot Core Ratio
Memory	Total / Free / Cached Memory, Total / Available Swap, Buffers
Disk	Disk I/O Queue Length, Read/Write Operation Count, Read/Write Latency, Average I/O Size
Process	Process Count, GPU-Using Process Count / CPU Utilization, Window-Focused Process CPU / Memory Utilization, Input Event
Scheduling	Task Migration Count, Lock Contention / Acquisition Failure Count, Lock Hold / Thread Blocked Time, Thread Wakeup / Context Switch Latency, Context Switch Count, Run Queue Length
Network	Total / New / Closed / Reset Connection Count, Data Sent/Received, Packet Count, Packet Retransmission Rate, TCP/UDP Ratio, Network Interrupt Count / Handling Latency, Data Transmit/Receive Latency, Kernel-Userspace Latency

Perception module collects data from multiple sources, including kernel-level eBPF programs, the procfs file system, and GNOME Shell for desktop environment information. This multi-faceted approach allows for a comprehensive view of system activity, from low-level kernel events to high-level user interactions, while maintaining minimal performance overhead. The feature collection process captures multiple dimensions of system behavior, as detailed in Table 1.

3.3 Decision module

The Decision module is the intelligent core of ASA, responsible for analyzing collected features and identifying workload patterns. This module employs a machine learning-based classification approach that has been trained on diverse workload scenarios to recognize common patterns such as CPU-intensive computing, I/O-bound operations, interactive applications, and mixed workloads.

The classification model uses an ensemble approach combining multiple algorithms to improve robustness and accuracy. The primary classifier is based on XGBoost, chosen for its excellent performance on structured data and ability to handle feature interactions. Additional classifiers including Random Forest and Support Vector Machines provide complementary perspectives and help detect edge cases.

To address the challenge of temporal stability in dynamic environments, ASA implements a Time-Weighted Probability Voting mechanism to aggregate classification results over a sliding window. This method ensures that scheduling decisions are based on stable workload patterns rather than transient fluctuations. The detailed procedure is formalized in Algorithm 1.

This approach effectively filters out transient system fluctuations while remaining responsive to genuine workload transitions.

The Policy Matching component maintains a “scenario-optimal scheduler” mapping table of as ASA’s long-term memory that maps workload patterns to best-performing

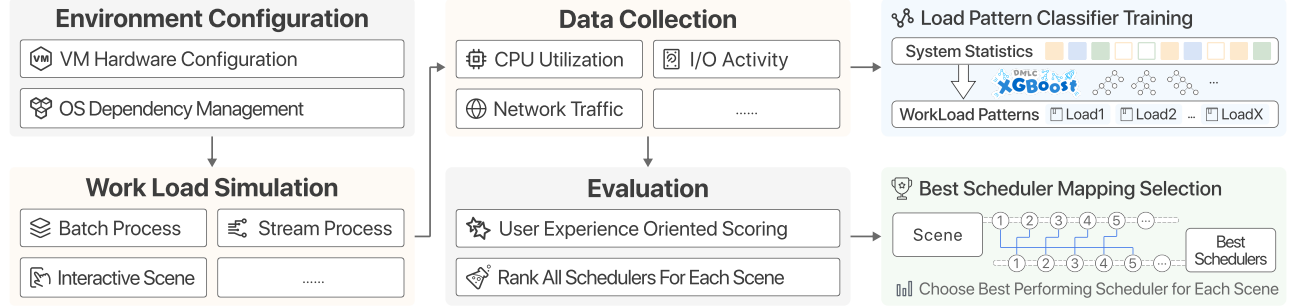


Figure 3. The Workflow within a Single Offline Preparation Stage.

Algorithm 1 Time-Weighted Probability Voting

Input: Sequence of probability distributions $\{P_1, \dots, P_W\}$ over a window W , threshold θ , decay factor α .

Output: Final selected class C_{final} .

```

1:
2: Initialize aggregated scores  $A_c = 0$  for all classes  $c$ .
3: for  $t = 1$  to  $W$  do
4:   for each class  $c$  with probability  $p_{t,c}$  in  $P_t$  do
5:     Let  $p'_{t,c} = \begin{cases} p_{t,c} & \text{if } p_{t,c} > \theta \\ 0 & \text{otherwise} \end{cases}$ 
6:      $A_c \leftarrow A_c + p'_{t,c} \cdot \alpha^{W-t}$ 
7:   end for
8: end for
9:  $C_{final} \leftarrow \arg \max_c A_c$ 
10: return  $C_{final}$ 

```

scheduling policies. The workload classifier model and this scheduler mapping table are generated through a specialized offline prepare pipeline, which will be discussed in detail in Section 4.

3.4 Action module

The Action module implements the execution phase of ASA’s decision-making process. Built on the Linux kernel’s sched_ext framework, this module enables dynamic switching between different scheduling policies without requiring system restart or interrupting running processes. The Action module continuously monitors the scheduler’s status. When the Decision module selects a new scheduling policy, the Action module checks if it is the current one. If not, it stops the current scheduler and switches to the new one. To prevent system instability and overhead from frequent switching, a cooldown period is enforced between policy changes.

The implementation leverages eBPF programs to implement custom scheduling logic that can be loaded and unloaded dynamically. This approach provides the flexibility needed for rapid policy switching while maintaining the safety and isolation guarantees of the kernel environment.

4 Offline Preparation

4.1 Cross-Platform Scheduling Challenge and Solution

The implementation of the Adaptive Scheduling Agent (ASA) must address the challenge of mapping dynamic workload patterns to optimal scheduling policies in heterogeneous hardware environments. This is complicated by two factors: the platform dependency of workload characteristics, where the same application exhibits different features on different hardware, and the hardware adaptability of scheduling policies, where the optimal policy for a given workload can vary with the hardware configuration. However, our core insight is that for any single machine, a workload pattern can correspond to an optimal scheduler. This principle enables a practical cross-platform optimization strategy, where the general task of pattern recognition can be decoupled from the machine-specific task of policy selection.

4.2 User Experience Oriented Evaluation Criteria

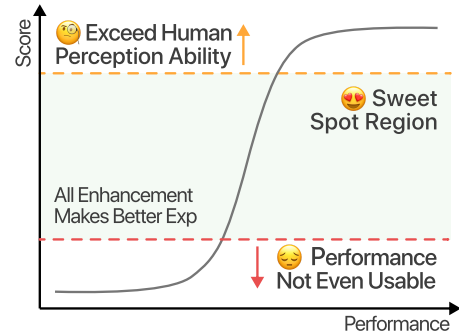


Figure 4. The performance evaluation curve for the User Experience Oriented Evaluation Criteria

Traditional scheduler evaluations often rely on system-level metrics like CPU utilization, which can have a semantic gap with the user’s perceived experience. To address this, we have developed a *User Experience Oriented Evaluation Criteria* that maps low-level scheduling behavior to high-level user perception, covering application response latency, UI

rendering smoothness, interactive feedback timeliness, and other user experience dimensions. This shifts the optimization goal from maximizing system efficiency to optimizing user experience.

The criteria work as a quantitative metric that mirrors the non-linear nature of human perception for evaluating the actual user experience. In the “Performance Not Even Usable” region (see Figure 4), variations in performance are largely meaningless, as the application is already functionally compromised—much like the indistinguishable user frustration (e.g., between 1 and 10 frames per second in a video game). Conversely, in the “Exceeds Human Perception” region, further enhancements offer diminishing returns, as they surpass the user’s perceptual limits on a given device (e.g., the difference between 200 and 300 FPS is often negligible). The most significant gains in user experience occur within the intervening ‘Sweet Spot Region,’ where users are highly sensitive to performance improvements. This also discourages trade-offs in multitask scenarios that sacrifice one task’s performance for marginal gains in another.

In our specific implementation, our user experience evaluation score is composed of two parts: *stability score* and *absolute value score*. The calculation formula for the stability score is:

$$S_{\text{stability}} = \frac{1}{1 + k_s \cdot \sigma^{\gamma_s}}$$

In this formula, σ is the data variance, and k_s and γ_s are the stability coefficient and exponent, respectively.

The absolute value score’s formula is constructed based on a modified sigmoid function:

$$S_{\text{value}}(x) = \frac{1}{1 + e^{\pm k_v \cdot (\frac{x-\theta}{f})^{\gamma_v}}}$$

In this formula, the synergistic effect of the value exponent γ_v and value coefficient k_v enables dynamic adjustment of the size of the “Sweet Spot Region”. The threshold θ defines the metric’s median or a critical performance point representing good experience, while the scaling factor f is selected via standard deviation or threshold to guarantee comparability of data with different units and magnitudes on the scoring scale. “ $\pm$ ” controls whether larger or smaller values are preferred. By integrating stability scoring and value scoring through weighted fusion, the final quantified score reflecting the user’s overall perceived experience in a specific scenario is obtained.

4.3 Model Training & Scheduler Mapping Selection

This part introduces a 3-stage strategy realizing a process of model updating from a coarse-grained level to a refined level.

4.3.1 Stage 1: Prototype Learning. The primary objective of this stage is to construct ASA’s foundational cognitive capabilities. This is achieved by training an initial workload

recognition model with generalized data collected in a standardized test environment and establishing a preliminary mapping between workload patterns and optimal schedulers.

To accomplish this, we first deploy different Linux schedulers to run completely through a series of pre-defined, typical test scenarios. During this process, two key types of data are collected simultaneously: 1) System Operational Metrics Data, as detailed in Table 1, which forms the feature vectors for training the workload classification model, and 2) Scenario Performance Metrics that will later be used to score and rank each scheduler.

With the system operational metrics, we train a preliminary workload classification model. Concurrently, based on the performance evaluation results, we determine the optimal scheduling strategy for each test scenario, thereby constructing an initial scheduler mapping table of “scenario-optimal scheduler”.

4.3.2 Stage 2: Dynamic Overhead Calibration. In real-world operations, scheduler switching is not a zero-cost action. To make the model predict accurately with the random perturbation of the cost of scheduler switching and system metrics monitoring, we perform a key modification to the scheduler set ASA uses: the schedulers now executes the complete pre-loading and validation procedures when invoked, but without actually being mounted as the system’s active scheduler, thus decoupling the overhead measurement process from the system’s current scheduling behavior.

We deploy ASA prototype based on the scheduler mapping table selected in the first stage. When ASA decides to switch schedulers based on the perceived system state, it invokes the modified “shadow” scheduler. This allows us to realistically record and simulate the composite overhead generated by system metric collection and scheduling decisions within the actual decision loop. The data collected in this stage is used to refine the workload classifier model and scheduler mapping table, ensuring that ASA can perceive the cost of each scheduler switch, thereby effectively preventing the frequent and ineffective decision oscillations caused by the perturbation of scheduler switching.

4.3.3 Stage 3: Generalization Model Training. To learn the profound impact of dynamic scheduler switching on an application’s runtime behavior, in this stage, we run ASA in fully functional mode which performs actual scheduler switching.

In this stage, ASA uses the calibrated mapping table from Stage 2 as its initial configuration to operate and autonomously perform scheduler switches in a live environment. Targeted comparative replacement tests are conducted for the ‘optimal schedulers’ in certain scenarios to explore better matching strategies and collect more system state data.

The collected state data is then used to further optimize the workload pattern classification model and the scheduler mapping table, enabling it to understand the causal impacts of different scheduling strategies on system behavior. Through this process, ASA’s model and mapping data are derived entirely from its own runtime environment. This significantly enhances its decision-making accuracy and generalization capability across diverse workloads and complex hardware environments.

4.4 Fast Adaptation to New Hardware

Deploying the generalized ASA agent onto a new hardware platform is a streamlined process. By exclusively running the “Generalization Model Training” (Stage 3) on the target machine, ASA can interact directly with the new environment and its available schedulers, resulting in the creation of a precise, hardware-specific scheduler mapping table ready for immediate use.

The entire process can then be followed by an optional, low-overhead fine-tuning process that leverages the newly collected system metrics to precisely align the model with the unique performance characteristics of the target machine for scenarios demanding maximum performance.

5 Evaluation

This section systematically evaluates the effectiveness of the Adaptive Scheduling Agent (ASA). We first detail the experimental methodology in Section 5.1. We then proceed to validate a series of claims central to our paper. In Section 5.2, we argue that ASA’s dynamic selection paradigm is fundamentally superior to mainstream single-scheduler approaches. In Section 5.3, we demonstrate that ASA’s performance effectively approaches that of an oracle in our scheduler set. In Section 5.4, we verify ASA’s ability to generalize robustly to new, unseen environments. Subsequently, Section 5.5 presents ablation studies to validate the efficacy of ASA’s key design components. Finally, Section 5.6 analyzes the system overhead to confirm its practicality, and Section 5.7 summarizes our findings and discusses limitations.

5.1 Experimental Setup

5.1.1 Evaluation Workloads and Platforms. To simulate realistic and challenging conditions, we constructed a benchmark suite of 28 scenarios. These are generated by pairing 4 interactive, latency-sensitive applications (Web Browsing, Audio Remix, Office File, Game Play) with 7 resource-intensive, background workloads (e.g., kernel compilation, blender render, LLM local generation). The detailed design of these scenarios can be found in Appendix A, B, and C.

This design creates resource contention and complex, mixed-workload environments, which are notoriously difficult for traditional schedulers and serve as a strong testbed for ASA’s dynamic decision-making.

Table 2. Benchmark Scenarios

No.	Scenario
S1	Game Play & Archive Extraction
S2	Game Play & Blender Render
S3	Game Play & Kernel Compile
S4	Game Play & LLM Generate
S5	Game Play & Disk IO
S6	Game Play & Network Transfer
S7	Game Play & Video Render
S8	Office File Editing & Archive Extraction
S9	Office File Editing & Blender Render
S10	Office File Editing & Kernel Compile
S11	Office File Editing & LLM Generate
S12	Office File Editing & Disk IO
S13	Office File Editing & Network Transfer
S14	Office File Editing & Video Render
S15	Web Browsing & Archive Extraction
S16	Web Browsing & Blender Render
S17	Web Browsing & Kernel Compile
S18	Web Browsing & LLM Generate
S19	Web Browsing & Disk IO
S20	Web Browsing & Network Transfer
S21	Web Browsing & Video Render
S22	Audio Remix & Archive Extraction
S23	Audio Remix & Blender Render
S24	Audio Remix & Kernel Compile
S25	Audio Remix & LLM Generate
S26	Audio Remix & Disk IO
S27	Audio Remix & Network Transfer
S28	Audio Remix & Video Render

Our evaluation is conducted on a cluster of 10 distinct virtual machines (VMs) managed by Proxmox VE, which allows us to adjust parameters such as the number of CPU cores allocated to a VM on the base hardware, thereby creating a richer set of test hardware platforms. As shown in Table 3, these VMs feature heterogeneous CPU architectures with core counts ranging from 2 to 20, including both symmetric core designs and asymmetric Intel Hybrid Architecture configurations. Memory capacity varies from 8GB to 16GB, while all VMs share a consistent, high-performance 64G SSD storage backend. This heterogeneity is intentional, designed to rigorously test ASA’s robustness and generalization capabilities. The entire process is managed by a bespoke automated system that standardizes benchmark execution, data collection, and analysis within isolated environments, guaranteeing credible and repeatable results. Within this cluster, four VMs (specifically, VM 120, 121, 125 and 128) are designated as **prototype VMs** for initial model calibration. The remaining

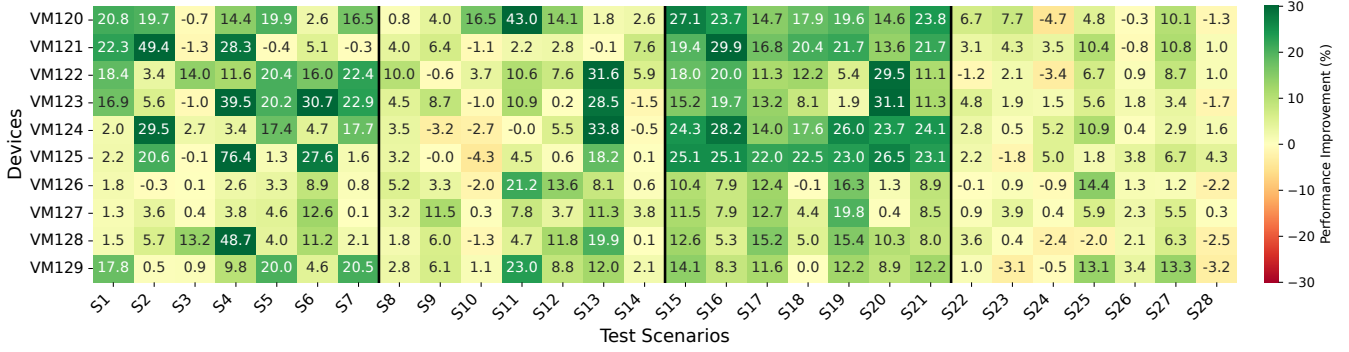


Figure 5. Global heatmap (scenarios × devices, relative to EEVDF)

six are treated as **unseen VMs** to evaluate how well ASA adapts to new environments.

Table 3. Evaluation Platform Configuration

VM ID	Base CPU	Cores	Threads	Mem
VM120	Intel i5-6500	4	4	8GB
VM121	Intel i5-6500	4	4	16GB
VM122	Intel i5-9400	6	6	12GB
VM123	Intel i5-9400	6	6	16GB
VM124	Intel i5-9400	4	4	8GB
VM125	Intel i5-9400	2	2	8GB
VM126	Intel i7-12700	8P+4E	20	16GB
VM127	Intel i7-12700	8P	16	16GB
VM128	Intel i7-12700	6P	12	16GB
VM129	Intel i7-12700	4P+4E	12	16GB

5.1.2 Baselines and Metrics. We evaluate ASA against two critical baselines:

- **EEVDF:** The default scheduler in Linux since kernel 6.6. As the latest evolution in general-purpose schedulers, it represents the mainstream in single-policy static optimization.
- **Static-Oracle:** A baseline representing the upper bound of performance achievable by any static scheduler within our expert set. For each of the 28 scenarios, the Static-Oracle score is the performance of the single best-performing expert scheduler, determined post-hoc.

Our primary metric is the *User Experience Oriented Criteria*, a normalized composite score previously detailed in Section 4. This approach bridges the gap between low-level system performance and tangible user satisfaction. We also report on relative performance rankings, including win/loss rates and Top-K performance rates, to provide a comprehensive picture. These metrics share a common comparison unit: a single test run of a specific scenario on a particular VM. The win rate is then the percentage of runs where ASA

surpasses the baseline, while the Top-K rate is the percentage of runs where ASA’s performance ranks among the top K of all available schedulers. The expert scheduler set available to ASA includes: *scx_p2dq* [9], *scx_bpfland* [2], *scx_nest* [24], *scx_lavd* [3], *scx_simple*, *scx_flash* [4], *scx_rusty* [1], and the baseline *EEVDF* [29].

5.2 Claim 1: Dynamic Selection Surpasses Static Optimization

A core premise of ASA is that the paradigm of dynamic policy selection is inherently more powerful than refining a single, static policy. Figure 5 and the win/loss statistics in Figure 6a substantiate this claim by comparing ASA’s performance against EEVDF. The results show that ASA achieves an overall win rate of 86.4% against EEVDF, with a global average improvement of +8.83% (95% CI [7.08%, 10.59%]). The distribution of these improvements is detailed in Figure 6c, which shows that the gains are not concentrated in a few outlier scenarios but are broadly distributed, with a clear positive skew.

This highlights a fundamental challenge for any single scheduler: a policy optimized for one type of workload (e.g., throughput for batch jobs) cannot guarantee optimal performance for another (e.g., latency for interactive tasks). EEVDF, despite its sophistication, cannot escape this trade-off. ASA overcomes this by switching to the most suitable expert policy on demand.

Notably, in many scenarios where ASA does not significantly outperform EEVDF, it is because EEVDF itself is the optimal or near-optimal choice for that specific workload. In these cases, ASA correctly identifies this and selects EEVDF (or a similarly performing scheduler), with its marginal overhead sometimes resulting in performance slightly below the native EEVDF. This is not a failure of ASA’s logic but a reflection of the current scheduler set’s boundaries. As the portfolio of expert schedulers expands, the frequency and magnitude of ASA’s performance gains are expected to increase further.

5.3 Claim 2: ASA Effectively Approaches Oracle Performance

A prerequisite for approaching oracle-level performance is the ability to accurately identify the current workload scenario to select the appropriate scheduler. The agent’s decision-making core—the workload classifier—achieves this with high precision. Even without any environment-specific fine-tuning, the base model correctly identifies the running workload scenario with 96.83% accuracy. This high baseline accuracy is fundamental to ASA’s ability to select an appropriate scheduler from the very beginning.

With this accurate workload identification as a foundation, we now evaluate how close ASA’s performance gets to the maximum achievable by its expert set. We measure this by comparing ASA to the Static-Oracle. The Top-K ranking statistics in Figure 6b are particularly revealing: ASA’s selected scheduler is the single best one in 45.4% of cases and ranks among the top three in 78.9% of cases. This demonstrates that ASA’s decision-making process is not just making a good choice, but is frequently making the best possible choice from the available options.

Furthermore, in a number of scenarios characterized by distinct phases or mixed workloads, ASA’s dynamic nature allows it to outperform the Static-Oracle. By switching policies mid-task to match changing workload characteristics, ASA can achieve a higher overall score than any single scheduler could. This “dynamic surpassing static” phenomenon underscores the unique value of ASA’s adaptive approach, which lies not in merely picking a global winner but in tracking the instantaneous optimum.

5.4 Claim 3: ASA Generalizes Effectively Across Diverse Environments

A practical intelligent system must be able to generalize from its training environment to new, unseen ones. To validate this, we compare ASA’s performance on the four prototype VMs against its performance on the six unseen VMs.

As shown in Figure 6d, ASA delivers consistent, positive performance improvements across all 10 VMs. The performance metrics on unseen machines are comparable to, and in some cases slightly better than, those on the prototype machines. Specifically, on prototype VMs, ASA achieves a win rate of 83.9% against EEVDF, with its choices ranking as Top-1 in 44.6% of scenarios and Top-3 in 76.8% of scenarios. On unseen VMs, these figures improve to a win rate of 88.1%, a Top-1 rate of 45.8%, and a Top-3 rate of 80.4%. This comparison reveals that performance on unseen machines is not just maintained but robustly transferred. This demonstrates that ASA’s three-stage calibration pipeline is effective at capturing the essential performance characteristics of a machine without overfitting to it, confirming its strong generalization capability and readiness for real-world deployment.

5.5 Ablation Studies

To isolate the impact of key design choices, we conducted ablation studies. A critical component of ASA is the Time-Weighted Probability Voting mechanism, designed to ensure decision stability by filtering out transient system noise. Figure 7 shows its impact by varying the window length W .

The results show a clear U-shaped curve for response delay. A very short window ($W < 4s$) leads to high average delay because the system becomes overly sensitive to transient noise, causing frequent, erroneous scheduler switches. This policy “thrashing” means the system spends significant time in a suboptimal state while oscillating, which paradoxically increases the time required to settle on the correct policy. Conversely, a very long window ($W > 10s$) also increases delay by making the system too slow to respond to genuine workload changes. The plot reveals a Pareto-optimal range between 4s and 7s. We chose $W = 6s$ as a robust compromise, as it significantly reduces both the error rate and the response delay compared to the extremes. This confirms that the voting mechanism is crucial for balancing responsiveness and stability, enabling ASA to make confident decisions.

We also evaluated the impact of the online fine-tuning mechanism. Without fine-tuning, the base workload classifier achieves a notable accuracy of 96.83%. After the online fine-tuning process, the accuracy of the raw model output (before the time-weighted voting is applied) increases to 99.19%. This demonstrates that the fine-tuning mechanism is highly effective at adapting the model to the specific nuances of a new environment, further enhancing the precision of ASA’s decision-making.

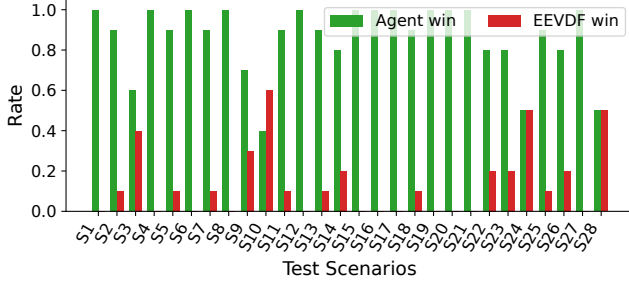
5.6 System Overhead

For ASA to be practical, its performance benefits must significantly outweigh its operational costs. We measured ASA’s resource consumption and latency across all test scenarios. On average, ASA’s user-space agent and kernel-level components collectively consume just 1.45% of a single CPU core and 297 MB of memory. This level of overhead is negligible on modern systems.

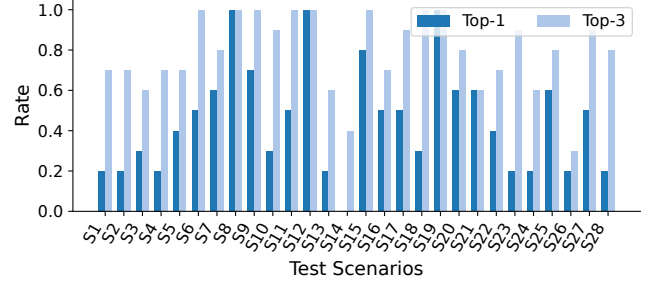
Table 4. ASA Latency Statistics (values in milliseconds)

Latency Type	Average	P90	P95	P99
Inference Latency	1.004	1.109	1.329	3.445
Decision Latency	17.305	0.083	0.307	376.417
Total Latency	18.309	1.359	2.722	379.540

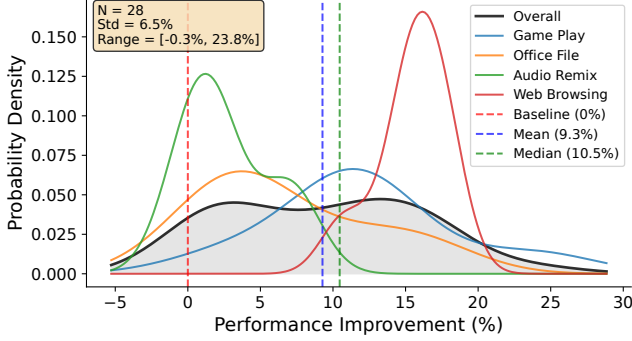
To assess the responsiveness of our agent, we measured the latency of its decision-making pipeline, breaking it down into two main components. *Inference latency* covers the time from data collection to the prediction of the workload pattern’s probability distribution. *Decision latency* is the time from the completion of inference to the determination of



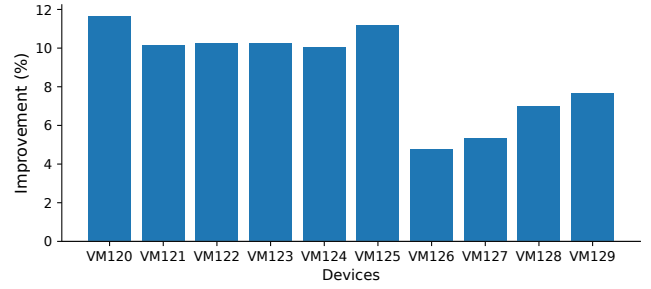
(a) Win/loss rate compared with EEVDF



(b) Top-K ranking within the scheduler set



(c) Performance improvement distribution density



(d) VM geometric mean performance improvement

Figure 6. Comprehensive ASA system performance evaluation: (a) Win/loss rate analysis reveals ASA’s advantage scenario distribution relative to baseline; (b) Top-K ranking statistics quantify ASA’s relative ranking performance across scenarios; (c) Performance improvement distribution density plot shows the overall distribution pattern and statistical characteristics of gains across all scenarios; (d) VM geometric mean performance improvements demonstrate consistent performance across different hardware configurations.

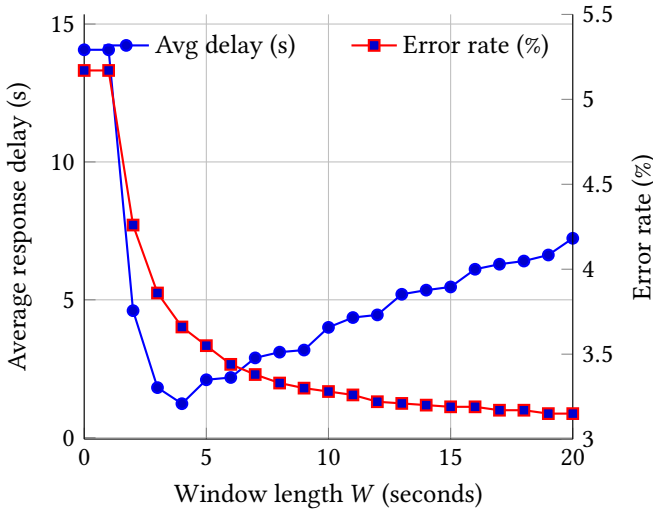


Figure 7. Sensitivity of Time-Weighted Probability Voting to window length

a scheduling policy switch. The total latency is the sum of these two.

Table 4 presents the average and percentile latencies for these components. To understand worst-case behavior, we examine the 90th, 95th, and 99th percentile (P90, P95, P99) latencies. These values indicate that 90%, 95%, and 99% of the decisions, respectively, are made faster than these thresholds. The average total latency is a mere 18.3 ms. Even in the worst-case scenarios, the P99 latency is 379.5 ms, demonstrating that ASA can react to system state changes in a timely manner even under heavy load. In summary, ASA’s benefits are delivered at a minimal and acceptable cost.

5.7 Discussion and Summary

Our comprehensive evaluation demonstrates that ASA successfully validates its core premise: dynamic, intelligent selection of expert schedulers is a superior paradigm for modern operating systems. The experiments confirm that ASA significantly outperforms a mainstream general-purpose scheduler (EEVDF), effectively approaches the performance of an oracle of our scheduler set, and generalizes robustly to new hardware environments. Ablation studies verify that its decision stability mechanisms are critical to its success, and overhead analysis confirms its real-world practicality.

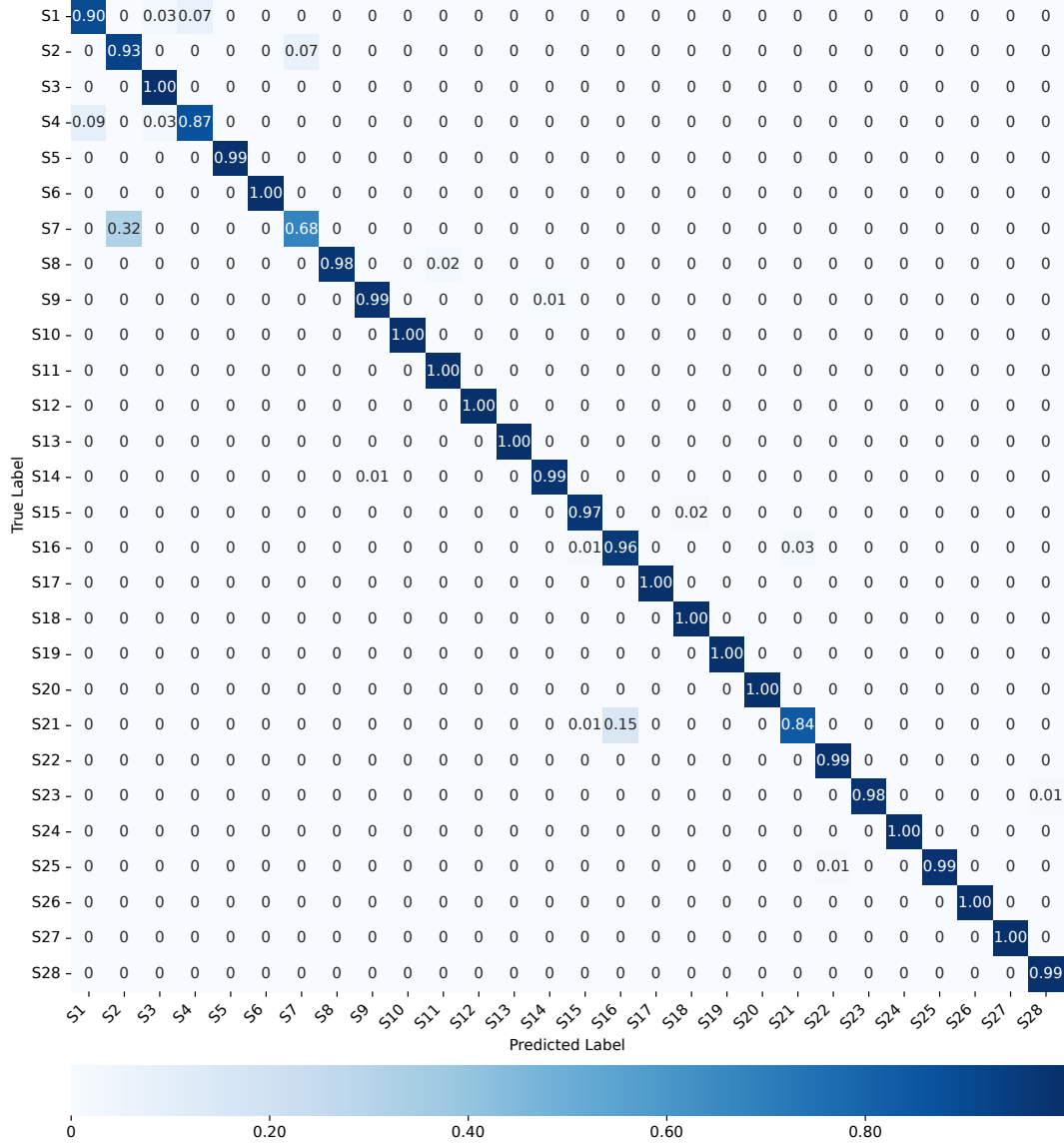


Figure 8. Workload Classifier Confusion Matrix

The primary limitation of ASA is that its performance ceiling is defined by the quality and diversity of its expert scheduler portfolio. However, this is also a strength. It provides a clear and scalable path for future improvement: as new, more advanced schedulers are developed by the community, they can be readily integrated into ASA’s framework, systematically raising the performance potential of the entire system.

6 Conclusion

By designing and implementing the Adaptive Scheduling Agent, this study has verified the feasibility and effectiveness of the “dynamically select the optimal policy” design philosophy in the field of operating system scheduling. ASA not

only achieves technical innovation but, more importantly, provides new ideas and methods for the intelligent development of operating system schedulers.

As computing environments continue to evolve and application workloads become increasingly complex, traditional static scheduling policies are hard to meet the performance demands of modern systems. The adaptive, intelligent scheduling direction represented by ASA provides a feasible technical path to address this challenge, and the results of this research not only have certain academic value but also provide useful references for practical applications in industry.

References

- [1] Righi Andrea, Heo Tejun, Marchevsky Dave, Vernet David, Hillion Jake, and likewhatevs. n.d.. `scx_rusty`: A Rust Crate for `scx_rusty` Algorithm. https://lib.rs/crates/scx_rusty. Accessed: 2025-08-20.
- [2] Righi Andrea, Heo Tejun, Vernet David, Hillion Jake, and likewhatevs. n.d.. `scx_bpfland`: A Rust Crate for `scx_bpfland` Algorithm. https://lib.rs/crates/scx_bpfland. Accessed: 2025-08-20.
- [3] Righi Andrea, Heo Tejun, Vernet David, Hillion Jake, and likewhatevs. n.d.. `scx_lavd`: A Rust Crate for `scx_lavd` Algorithm. https://lib.rs/crates/scx_lavd. Accessed: 2025-08-20.
- [4] Righi Andrea, Heo Tejun, and Daniel Hodges. 2025. `scx_flash`: A Novel Scheduling Algorithm for Real-Time Systems. *Journal of Embedded Systems* 10, 3 (2025), 123–145.
- [5] Ansen88. 2025. *Hackbench*. GitHub. <https://github.com/Ansen88/hackbenchv2.3>
- [6] Esmail Asyabi, Azer Bestavros, Renato Mancuso, Richard West, and Erfan Sharafzadeh. 2020. Akita: A cpu scheduler for virtualized clouds. *arXiv preprint arXiv:2009.09104* (2020).
- [7] Blender Foundation. 2025. Blender - a 3D modelling and rendering package. <https://www.blender.org>
- [8] Li Chen, Justinas Lingys, Kai Chen, and Feng Liu. 2018. AuTO: scaling deep reinforcement learning for datacenter-scale automatic traffic optimization. In *Proceedings of the 2018 Conference of the ACM Special Interest Group on Data Communication (Budapest, Hungary) (SIGCOMM '18)*. Association for Computing Machinery, New York, NY, USA, 191–205. doi:10.1145/3230543.3230551
- [9] Hodges Daniel, Heo Tejun, and Hillion Jake. n.d.. `scx_p2dq`: A Rust Crate for `scx_p2dq` Algorithm. https://lib.rs/crates/scx_p2dq. Accessed: 2025-08-20.
- [10] DeepSeek-AI, Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, Xiaokang Zhang, Xingkai Yu, Yu Wu, Z. F. Wu, Zhibin Gou, Zhihong Shao, Zhuoshu Li, Ziyi Gao, Aixin Liu, Bing Xue, Bingxuan Wang, Bochao Wu, Bei Feng, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Qu, Hui Li, Jianzhong Guo, Jiashi Li, Jiawei Wang, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, J. L. Cai, Jiaqi Ni, Jian Liang, Jin Chen, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Liang Zhao, Litong Wang, Liyue Zhang, Lei Xu, Leyi Xia, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Meng Li, Miaojun Wang, Mingming Li, Ning Tian, Panpan Huang, Peng Zhang, Qiancheng Wang, Qinyu Chen, Qiushi Du, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, R. J. Chen, R. L. Jin, Ruyi Chen, Shanghao Lu, Shangyan Zhou, Shanhuan Chen, Shengfeng Ye, Shiyu Wang, Shuiping Yu, Shunfeng Zhou, Shutong Pan, S. S. Li, Shuang Zhou, Shaoqing Wu, Shengfeng Ye, Tao Yun, Tian Pei, Tianyu Sun, T. Wang, Wangding Zeng, Wanbiao Zhao, Wen Liu, Wenfeng Liang, Wenjun Gao, Wenqin Yu, Wentao Zhang, W. L. Xiao, Wei An, Xiaodong Liu, Xiaohan Wang, Xiaokang Chen, Xiaotao Nie, Xin Cheng, Xin Liu, Xin Xie, Xingchao Liu, Xinyu Yang, Xinyuan Li, Xuecheng Su, Xuheng Lin, X. Q. Li, Xiangyue Jin, Xiaojin Shen, Xiaosha Chen, Xiaowen Sun, Xiaoxiang Wang, Xinnan Song, Xinyi Zhou, Xianzu Wang, Xinxia Shan, Y. K. Li, Y. Q. Wang, Y. X. Wei, Yang Zhang, Yanhong Xu, Yao Li, Yao Zhao, Yaofeng Sun, Yaohui Wang, Yi Yu, Yichao Zhang, Yifan Shi, Yiliang Xiong, Ying He, Yishi Piao, Yisong Wang, Yixuan Tan, Yiyang Ma, Yiyuan Liu, Yongqiang Guo, Yuan Ou, Yuduan Wang, Yue Gong, Yuheng Zou, Yujia He, Yunfan Xiong, Yuxiang Luo, Yuxiang You, Yuxuan Liu, Yuyang Zhou, Y. X. Zhu, Yanhong Xu, Yanping Huang, Yaohui Li, Yi Zheng, Yuchen Zhu, Yunxian Ma, Ying Tang, Yukun Zha, Yuting Yan, Z. Z. Ren, Zehui Ren, Zhangli Sha, Zhe Fu, Zhean Xu, Zhenda Xie, Zhengyan Zhang, Zhewen Hao, Zhicheng Ma, Zhigang Yan, Zhiyu Wu, Zihui Gu, Zijia Zhu, Zijun Liu, Zilin Li, Ziwei Xie, Ziyang Song, Zizheng Pan, Zhen Huang, Zhipeng Xu, Zhongyu Zhang, and Zhen Zhang. 2025. DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning. arXiv:2501.12948 [cs.CL] <https://arxiv.org/abs/2501.12948>
- [11] Christina Delimitrou and Christos Kozyrakis. 2013. Paragon: QoS-aware scheduling for heterogeneous datacenters. *ACM SIGPLAN Notices* 48, 4 (2013), 77–88.
- [12] Christina Delimitrou and Christos Kozyrakis. 2014. Quasar: Resource-efficient and qos-aware cluster management. *ACM Sigplan Notices* 49, 4 (2014), 127–144.
- [13] Allen B Downey. 1997. *A model for speedup of parallel programs*. University of California, Berkeley, Computer Science Division.
- [14] ESnet. 2025. iperf3: A TCP, UDP, and SCTP network bandwidth measurement tool. <https://software.es.net/iperf/>.
- [15] FabricMC. 2024. Fabric API. <https://fabricmc.net/>.
- [16] Michelle Galin and Anna Hedblom. 2025. *Linux Kernel Scheduler Evaluation for Performance-Critical Telecom Workloads*. Master's thesis. Linköping University, Department of Computer and Information Science.
- [17] Georgi Gerganov. 2024. llama.cpp. <https://github.com/ggerganov/llama.cpp>.
- [18] Philippe Gerum. 2004. Xenomai-Implementing a RTOS emulation framework on GNU/Linux. *White Paper, Xenomai* 81 (2004).
- [19] Jack Tigar Humphries, Neel Nattu, Ashwin Chaugule, Ofir Weisse, Barret Rhoden, Josh Don, Luigi Rizzo, Oleg Rombakh, Paul Jack Turner, and Christos Kozyrakis. 2021. ghOST: Fast and Flexible User-Space Delegation of Linux Scheduling. In *Proceedings of the ACM SIGOPS 28th Symposium on Operating Systems Principles CD-ROM*. New York, NY, USA, 588–604. <https://doi.org/10.1145/3477132.3483542>
- [20] Weiwei Jia, Jianchen Shan, Tsz On Li, Xiaowei Shang, Heming Cui, and Xiaoning Ding. 2020. vSMT-IO: Improving I/O Performance and Efficiency on SMT Processors in Virtualized Clouds. In *2020 USENIX Annual Technical Conference (USENIX ATC 20)*. 449–463.
- [21] Yuekai Jia, Kaifu Tian, Yuyang You, Yu Chen, and Kang Chen. 2024. Skyloft: A General High-Efficient Scheduling Framework in User Space. In *Proceedings of the ACM SIGOPS 30th Symposium on Operating Systems Principles (Austin, TX, USA) (SOSP '24)*. Association for Computing Machinery, New York, NY, USA, 265–279. doi:10.1145/3694715.3695973
- [22] Kostis Kaffes, Timothy Chong, Jack Tigar Humphries, Adam Belay, David Mazières, and Christos Kozyrakis. 2019. Shinjuku: Preemptive Scheduling for μ second-scale Tail Latency. In *16th USENIX Symposium on Networked Systems Design and Implementation (NSDI 19)*. 345–360.
- [23] Marios Kogias, George Prekas, Adrien Ghosn, Jonas Fietz, and Edouard Bugnion. 2019. R2P2: Making RPCs first-class datacenter citizens. In *2019 USENIX Annual Technical Conference (USENIX ATC 19)*. 863–880.
- [24] Julia Lawall, Himadri Chhaya-Shailesh, Jean-Pierre Lozi, Baptiste Lepers, Willy Zwaenepoel, and Gilles Muller. 2022. OS scheduling with nest: keeping tasks close together on warm cores. In *Proceedings of the Seventeenth European Conference on Computer Systems*. 368–383.
- [25] Chung Laung Liu and James W Layland. 1973. Scheduling algorithms for multiprogramming in a hard-real-time environment. *Journal of the ACM (JACM)* 20, 1 (1973), 46–61.
- [26] Weihong Liu, Jiawei Geng, Zongwei Zhu, Jing Cao, and Zirui Lian. 2022. Sniper: cloud-edge collaborative inference scheduling with neural network similarity modeling. In *Proceedings of the 59th ACM/IEEE Design Automation Conference (San Francisco, California) (DAC '22)*. Association for Computing Machinery, New York, NY, USA, 505–510. doi:10.1145/3489517.3530474
- [27] David Lo, Liqun Cheng, Rama Govindaraju, Parthasarathy Ranganathan, and Christos Kozyrakis. 2015. Heracles: Improving resource efficiency at scale. In *Proceedings of the 42nd Annual International Symposium on Computer Architecture*. 450–462.

- [28] Robert Love. 2005. *Linux Kernel Development* (2nd ed.). Novell Press, Chapter 4. This chapter, "The Process Scheduler", provides a canonical description of the O(1) scheduler's design and implementation as it appeared in the Linux 2.6 kernel..
- [29] LWN.net. n.d.. Earliest Eligible Virtual Deadline First (EEVDF) scheduler. Online. <https://lwn.net/Articles/925371/>
- [30] Hongzi Mao, Mohammad Alizadeh, Ishai Menache, and Srikanth Kandula. 2016. Resource management with deep reinforcement learning. In *Proceedings of the 15th ACM workshop on hot topics in networks*. 50–56.
- [31] Hongzi Mao, Malte Schwarzkopf, Shaileshh Bojja Venkatakrishnan, Zili Meng, and Mohammad Alizadeh. 2019. Learning scheduling algorithms for data processing clusters. In *Proceedings of the ACM special interest group on data communication*. 270–288.
- [32] Mojang Studios. 2011. Minecraft. Stockholm, Sweden. Java Edition.
- [33] Ingo Molnár. 2007. [patch] Modular Scheduler Core and Completely Fair Scheduler [CFS] - core bits. Linux Kernel Mailing List (LKML). <https://lkml.org/lkml/2007/4/13/180> The announcement and cover letter for the Completely Fair Scheduler patch series..
- [34] Igor Pavlov. 2025. 7-Zip. <https://www.7-zip.org/>.
- [35] George Prekas, Marios Kogias, and Edouard Bugnion. 2017. Zygus: Achieving low tail latency for microsecond-scale networked tasks. In *Proceedings of the 26th Symposium on Operating Systems Principles*. 325–341.
- [36] Tang Qinan, Gao Xing, Li Guilin, and Lin Juncong. 2024. Optimizing Linux Scheduling Based on Global Runqueue with SCX. In *2024 IEEE International Conference on Systems, Man, and Cybernetics (SMC)*. 1813–1819. doi:10.1109/SMC54092.2024.10831230
- [37] Qwen, An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. 2025. Qwen2.5 Technical Report. arXiv:2412.15115 [cs.CL] <https://arxiv.org/abs/2412.15115>
- [38] Federico Reghenzani, Giuseppe Massari, and William Fornaciari. 2019. The real-time linux kernel: A survey on preempt_rt. *ACM Computing Surveys (CSUR)* 52, 1 (2019), 1–36.
- [39] Pablo Sanabria, Tomás Felipe Tapia, Rodrigo Toro Icarte, and Andres Neyem. 2022. Solving Task Scheduling Problems in Dew Computing via Deep Reinforcement Learning. *Applied Sciences* 12, 14 (2022). doi:10.3390/app12147137
- [40] Selenium Project. 2024. Selenium WebDriver. <https://www.selenium.dev/>.
- [41] Xiaoxuan Sun, Yifei Duan, Yingnan Deng, Fan Guo, Guohui Cai, and Yuting Peng. 2025. Dynamic Operating System Scheduling Using Double DQN: A Reinforcement Learning Approach to Task Optimization. *arXiv preprint arXiv:2503.23659* (2025).
- [42] The Chromium Authors. 2024. Chromium. <https://www.chromium.org/>.
- [43] The Document Foundation. 2024. LibreOffice. <https://www.libreoffice.org/>
- [44] The Document Foundation. 2024. UNO - Universal Network Objects. In LibreOffice API Documentation. <https://api.libreoffice.org/>
- [45] The FFmpeg Project. 2025. FFmpeg. <https://ffmpeg.org/>.
- [46] The PulseAudio Developers. 2024. PulseAudio. <https://www.freedesktop.org/wiki/Software/PulseAudio/>.
- [47] Linus Torvalds and the Linux kernel community. 2025. The Linux Kernel. <https://www.kernel.org/>.
- [48] David Vernet. 2024. The current status and future potential of sched_ext. In *Linux Plumbers Conference 2024* (Vienna, Austria). <https://lpc.events/event/18/contributions/1694/>
- [49] Suchuan Xing, Yihan Wang, and Wenhe Liu. 2025. Self-Adapting CPU Scheduling for Mixed Database Workloads via Hierarchical Deep Reinforcement Learning. *Symmetry* 17, 7 (2025). doi:10.3390/sym17071109
- [50] Hang Zhu, Kostis Kaffes, Zixu Chen, Zhenming Liu, Christos Kozyrakis, Ion Stoica, and Xin Jin. 2020. RackSched: A Microsecond-Scale scheduler for Rack-Scale computers. In *14th USENIX Symposium on Operating Systems Design and Implementation (OSDI 20)*. 1225–1240.

A Interactive Scenario Design

Web Browsing Scenario: Users are particularly sensitive to issues such as page load latency, scrolling stutter, and unresponsive input. The scheduler's performance in this scenario directly affects the user's perceived experience and places high demands on the scheduler's ability to manage priorities and control latency for compute- or I/O-intensive tasks.

To objectively quantify the user experience in the web browsing scenario, this study designed the following test method: multiple local HTML files containing common interface components found in web browsing were written as test content. The Selenium framework [40], widely used for browser automation testing, was utilized to perform automated operations in the Chrome browser [42]. By injecting JavaScript timers, the response time of interactions such as clicks, scrolls, and tab switches, as well as the stability of the frame rate, were accurately recorded to evaluate the scheduler's performance in this scenario.

Game Play Scenario: Compared to other interactive scenarios, the Game Play scenario has significantly more intense content changes and user input. If the user's actions are not responded to immediately or the user cannot see the response in time, the experience will be greatly diminished. The Game Play scenario places extremely high demands on real-time performance and the consistency of frame generation time, thus it can also well reflect the scheduler's ability to schedule tasks with strict deadlines and its resource isolation capabilities.

This study selected the best-selling game of all time, "Minecraft" [32], for testing. As a sandbox game, it not only needs to continuously render a complex user perspective, but its built-in server also needs to perform a large amount of world simulation calculations (such as physics, creature behavior, chunk loading, etc.), which can comprehensively reflect the characteristics of various computing tasks under a gaming load. This study used the Fabric API [15] to develop a dedicated performance testing component to continuously simulate player movement from a fixed coordinate point on a fixed map, while collecting frame rate, frame generation time, number of loaded chunks, and MSPT (milliseconds per game tick) to quantify the scheduler's performance in handling high game loads and real-time requirements.

Audio Remix Scenario: When playing audio, the kernel needs to continuously fill the corresponding buffer of the sound card with audio data. If the buffer underflows due

to untimely responses, it can produce noticeable pops or interruptions, leading to a poor user experience. In a multi-track mixing Audio Remix scenario, the audio being played also needs to undergo complex real-time calculations and effects processing before output, which places extremely high demands on the scheduler’s latency control and real-time task assurance capabilities.

To quantify the performance in the Audio Remix scenario, this study used the PulseAudio suite [46] to perform multi-track audio mixing and effects application, and recorded data such as the number of buffer underruns and computation latency during the audio preview process to evaluate the scheduler’s ability to meet strict audio timing requirements.

Office File Editing Scenario: In document processing, users often perform a large number of frequent but brief interactive operations (such as text input, menu clicks, format adjustments) along with periodic or bursty background tasks (such as auto-save, spell check, etc.). The generation time of the related computational load is often unpredictable, while users have high demands for immediate feedback on their operations. This places high demands on the scheduler to balance foreground and background latency requirements and resource utilization.

This study chose the full-featured LibreOffice suite [43] for testing. By using its automation framework UNO [44], it simulated typical user operations such as content editing, format modification, and cross-document data copying for text documents, spreadsheets, and presentations. Key performance data such as input latency and rendering latency were recorded to evaluate the scheduler’s ability to balance user experience response speed and system resource utilization under this mixed workload.

B Streaming Computation Scenario Design

Disk IO Scenario: Due to the significant differences in speed and latency between storage media and memory/registers, programs often need to block and wait for a period of time when performing I/O operations, and are awakened after the requested operation is completed. The diverse latency characteristics and operational costs of Disk IO operations further increase the complexity for the scheduler to wake up processes, placing high demands on the scheduler’s ability to manage and schedule tasks that depend on different I/O resources.

To simulate the Disk IO scenario and quantify scheduler performance, the Disk IO scenario test creates a large number of data files ranging in size from KB to GB. At the same time, multiple processes are created to perform sequential or random read/write operations on randomly selected files and random locations within them. This simulates a highly complex concurrent I/O load. Key metrics such as IOPS (I/O operations per second), sequential read/write rates, and random read/write rates are recorded to measure the scheduler’s

ability to guarantee system throughput and response efficiency under heavy I/O pressure.

LLM Generate Scenario: With the rapid development of LLMs, the performance of low-parameter models has gradually reached a practical level, and the demand for local deployment is growing. When users interact with a locally deployed LLM, they usually input a prompt at any time and expect the model to respond quickly and complete the generation as soon as possible. This bursty and sustained compute-intensive load places extremely high demands on the scheduler’s timeliness in responding to user requests and its ability to schedule and manage large amounts of computational resources.

Considering the hardware resource conditions of the tested devices, this study chose to use a 1.5B model distilled from DeepSeek-R1 [10] to Qwen2.5 [37] for testing, and used llama-cpp [17] as the inference engine. It simulated user input of prompts to trigger model generation and recorded key metrics such as time-to-first-token and tokens-per-second generation speed to comprehensively evaluate the scheduler’s performance under this new type of computational load.

Network Transfer Scenario: Network transmission is an important component for the normal operation of many user applications. Users are very sensitive to the transmission rate, stability, and latency of the network when performing operations such as streaming media loading and online gaming. At the same time, the sending and receiving of network data packets involve frequent device I/O waits and interrupt handling. The large number of network events (interrupts, data copies, etc.) generated by high bandwidth utilization will bring a significant load to the CPU. These characteristics require the scheduler to be able to respond to network events in a timely manner, efficiently manage interrupt loads, and provide the necessary computing resources for network processes in time when handling network-related tasks. This places high demands on the scheduler’s scheduling strategy when handling high-concurrency network I/O and computing tasks.

To quantify the scheduler’s performance in the Network Transfer scenario, this study used the locally deployed network performance testing tool iperf3 [14] for testing. It simulated a high-throughput network data transmission load (TCP/UDP) and recorded key metrics such as data transmission rate, number of packets, and packet loss rate to reflect the scheduler’s ability to guarantee system network throughput, transmission efficiency, and stability when handling a large number of network I/O interrupts and data processing tasks.

C Batch Processing Scenario Design

Video Render Scenario: Video rendering, which simulates the export or transcoding process after professional

users edit a video, is a typical compute-intensive and time-consuming background task. Although not a direct user interaction, its completion time and the impact on system responsiveness when the user is performing other operations in the foreground are crucial to the overall user experience.

To quantify the scheduler's performance in this scenario, this study utilized FFmpeg [45], a widely used multimedia processing tool, for testing. A standard video source and preset output parameters were selected to simulate an actual video export or transcoding task. The total rendering completion time and average rendering rate were recorded to evaluate the scheduler's resource management efficiency under a sustained compute-intensive background load and its ability to guarantee the foreground user experience.

Blender Render Scenario: The Blender Render scenario can typically utilize multi-core processor resources in a highly parallel manner and consumes a huge amount of CPU and memory. It is a typical compute-intensive batch processing load. It can evaluate how effectively the scheduler schedules and manages a large number of parallel computing tasks to maximize the rendering throughput of a multi-core system. It also considers the system's responsiveness to bursty foreground interactive requests during high-intensity background rendering.

To quantify the scheduler's performance in this scenario, this study used Blender [7] to render a preset complex 3D scene using all available CPU cores as the test load. The rendering speed of tiles and the total rendering time were recorded to evaluate the scheduler's efficiency and multi-core resource utilization when dealing with highly parallel computing loads, and its impact on the responsiveness of basic foreground operations was observed.

Kernel Compile Scenario: The Kernel Compile scenario simulates the compilation process of an operating system kernel. It is a typical batch processing workload with extremely complex dependencies, and a dense mix of I/O read/write and computing tasks. The compilation process generates a large number of processes and threads and frequently switches between file reading/writing and computation. It can evaluate the scheduler's efficiency in handling such large-scale, mixed-type parallel tasks and whether the system responsiveness during high-intensity background compilation is sufficient to allow the user to perform basic interactive operations.

To quantify the scheduler's performance in this scenario, this study chose the complete compilation of the Linux 6.14 kernel [47] under the default configuration as the test load. The total compilation completion time was recorded to measure whether the scheduler can efficiently complete the background compute-intensive task while ensuring the smooth operation of foreground applications.

Compression/Decompression Scenario: The compression/decompression scenario simulates the sustained file

system access and/or computational load during the decompression or compression of large compressed files. It can evaluate how effectively the scheduler manages and schedules this type of task, which is highly dependent on or has a mixed dependency on disk I/O performance, and whether the responsiveness of foreground user applications is significantly affected during large-scale background compression/decompression operations.

To quantify the scheduler's performance in this scenario, this study used 7-Zip [34] to test the compression and decompression processes with different dictionary sizes. The compression/decompression rate was recorded to measure the scheduler's efficiency in handling large amounts of data and potentially mixed loads.