

Optimization of Cloud Computing Resource Allocation Through Artificial Intelligence-Based Scheduling Algorithms

Nesan Ryan

Department of Computer Science and Engineering, University of Nevada, Reno, Reno, NV, USA.

ryan384@unr.edu

Scott L. Taylor

Department of Computer Science, University of New Hampshire, Durham, NH, USA.

helloscott@unh.edu

Derran Anderson

Department of Computer Science, University of Alabama at Birmingham, Birmingham, AL, USA.

darren1994@uab.edu

Abstract

Cloud computing has become a foundational infrastructure for modern digital services, yet the allocation of its computational, storage, and network resources remains a persistent system-level challenge. Conventional scheduling policies often rely on static thresholds, heuristic placement rules, and reactive adjustments that struggle to handle heterogeneous workloads, rapid demand fluctuations, and long-term operational objectives. Artificial intelligence-based scheduling algorithms have emerged as a promising alternative because they can learn workload patterns, infer performance interference, and generate allocation decisions that are difficult to express through hand-crafted policies. This paper presents a system-oriented analysis of AI-driven resource allocation rather than a narrowly algorithmic contribution. It examines the architectural context of cloud scheduling, the main families of AI-based scheduling methods, structural trade-offs between centralized and decentralized control, and the governance, fairness, robustness, deployment, sustainability, and policy dimensions that shape real-world adoption. The discussion draws on simulation frameworks, production cluster management systems, and foundational learning methods to identify where AI-based schedulers provide leverage and where they introduce new operational risks. The paper argues that effective AI-based resource allocation requires not only accurate predictive models and stable learning algorithms, but also careful integration with organizational controls, transparency mechanisms, energy management, and accountability structures. Future research should address robustness under data drift, explainability for human operators, cross-layer coordination, and the sociotechnical conditions under which learned scheduling policies can be safely deployed in large-scale cloud infrastructures.

Keywords

cloud computing, resource allocation, artificial intelligence, scheduling algorithms, reinforcement learning, governance, robustness, sustainability, fairness.

1. Introduction

Cloud computing has been formalized as a model for enabling ubiquitous, convenient, on-demand network access to a shared pool of configurable computing resources that can be rapidly provisioned and released with minimal management effort [1]. This utility-oriented vision has shaped expectations about elastic capacity, multitenancy, and pay-per-use economics [2]. Economic and technical analyses of cloud platforms have further emphasized that elastic resource provisioning can reduce capital expenditure and improve utilization, but only when allocation decisions are sufficiently responsive to changing workload conditions [3]. At the same time, resource management in cloud environments has grown more complex because workloads are heterogeneous, performance interference is common, and operator goals are frequently multidimensional [4]. In this context, traditional scheduling heuristics often produce acceptable average behavior but fail to adapt to workload shifts, service-level objective violations, or infrastructure degradation. A general formulation of sequential decision making under uncertainty provides the conceptual foundation for learning-based schedulers that improve their policies from operational experience [5].

The motivation for applying artificial intelligence to cloud resource allocation is not simply to replace existing schedulers with more complex models. Instead, the central opportunity is to create allocation systems that can reason over large volumes of operational telemetry, capture nonlinear interactions among colocated workloads, and balance competing objectives such as performance, cost, energy consumption, and fairness. Recent demonstrations of deep reinforcement learning for resource management indicate that such methods can learn useful policies directly from system state signals [6]. However, translating research prototypes into production infrastructure raises broader questions about architectural fit, control authority, auditability, and resilience. These concerns are system-level rather than purely algorithmic. They involve the structure of the resource management plane, the granularity of decision intervals, the quality of monitoring data, the degree of human oversight, and the mechanisms through which learned decisions can be reviewed or overridden.

This paper provides an interdisciplinary analysis of AI-based scheduling algorithms for cloud resource allocation. It does not propose a single optimal algorithm or a mathematical formulation. Rather, it develops a conceptual framework for understanding how AI techniques interact with cloud architectures, governance arrangements, deployment constraints, sustainability goals, and policy requirements. The analysis emphasizes structural trade-offs, including centralization versus decentralization, model complexity versus interpretability, short-term responsiveness versus long-term stability, and efficiency versus fairness. The discussion draws on foundational cloud computing research, simulation platforms, production cluster managers, and learning-based systems to illustrate both the potential and the limits of AI-driven allocation. The paper concludes by outlining a research agenda oriented toward robust, accountable, and sustainable AI-based resource management in large-scale cloud infrastructures.

2. Foundations of Cloud Resource Allocation and Scheduling

Cloud resource allocation operates within a layered architecture that spans physical hosts, virtualization or containerization layers, cluster management services, and application workloads. The resource manager must decide where to place tasks, how many replicas to maintain, when to migrate workloads, and how to partition shared capacities across competing services. These decisions are complicated by the diversity of cloud workloads, which range from latency-sensitive interactive services to throughput-oriented batch jobs and data analytics pipelines. A comprehensive survey of cloud resource management identifies the

core challenges as multidimensional optimization, dynamic workload behavior, performance isolation, and uncertainty in resource demand [7]. Simulation toolkits have become important for evaluating allocation policies because they allow controlled experimentation with different workload mixtures, failure models, and scheduling configurations without risking production infrastructure [8].

Energy-aware dynamic consolidation has been a central concern in cloud resource allocation because data center energy consumption depends heavily on host utilization and the placement of virtual machines. Research on consolidation has shown that intelligent migration and packing policies can reduce energy consumption while maintaining service-level agreements [9]. Dynamic resource allocation using virtual machines has similarly demonstrated that adjusting allocations in response to monitored demand can improve resource utilization and application performance [10]. These approaches, however, often rely on threshold-based rules or greedy heuristics that assume relatively stable workload behavior. When demand becomes bursty, when applications exhibit complex resource interdependencies, or when performance interference increases, such policies require frequent manual tuning. This limitation has motivated the shift toward learning-based methods that can derive allocation policies from historical and real-time telemetry.

The operational context of cloud scheduling also includes the distinction between online and offline decisions. Online schedulers must make placement and scaling choices as requests arrive, while offline analyzers can periodically recompute plans for longer horizons. AI-based scheduling can be introduced at either level, but each level has different latency, data freshness, and risk constraints. For example, a real-time container placement decision may require millisecond-level inference, whereas a periodic capacity planner can afford more complex model evaluation. These architectural differences matter because they determine which AI techniques are feasible, how much contextual information can be included, and how failures can be contained. A robust AI-based allocation system therefore must be designed not as an isolated component but as part of a broader management plane with clear interfaces to monitoring, configuration, and incident response systems.

3. Artificial Intelligence Scheduling Paradigms

Artificial intelligence approaches to resource scheduling can be broadly grouped into supervised learning for prediction, reinforcement learning for sequential decision making, and deep learning for representation extraction. Supervised learning methods are often used to forecast workload demand, classify workload types, or predict the performance impact of colocation. Scalable tree boosting systems have been applied successfully to tabular operational data because they provide strong predictive accuracy and can handle mixed feature types [11]. Deep learning methods have expanded the ability to model complex, high-dimensional patterns in system telemetry and workload traces [12]. These predictive models can inform schedulers by converting uncertain future demand into expected resource requirements, but they do not by themselves generate allocation actions.

Reinforcement learning offers a more direct framework for learning scheduling policies through interaction with the environment. The success of deep reinforcement learning in complex sequential decision tasks has encouraged its application to resource management problems [13]. Unlike supervised models that predict a target value, reinforcement learning agents learn to select actions that maximize cumulative reward over time. In cloud scheduling, the reward might combine throughput, latency, energy consumption, and violation penalties. This makes reinforcement learning attractive for allocation tasks that involve delayed

consequences, such as the long-term effect of placing a workload on a host that will later experience thermal or power constraints. Attention mechanisms have further enabled models to weigh relevant portions of long input sequences, which is useful for learning workload patterns over time and for representing dependency structures in job graphs [14]. Reinforcement learning has also been explored for device placement optimization, where the learned policy assigns computational operations to heterogeneous hardware resources [15].

Despite these advances, applying AI scheduling paradigms in cloud systems requires more than model selection. Supervised predictors can be well calibrated on historical data but may fail when workload distributions shift. Reinforcement learning agents can exploit simulator inaccuracies or produce policies that are difficult to interpret. Deep learning models may require substantial training data and computational resources before they provide reliable guidance. Therefore, AI-based scheduling systems often combine multiple paradigms. For example, a workload predictor may generate inputs for a model predictive controller, or a reinforcement learning agent may be constrained by rule-based safety checks. The system-level design must address how these components interact, how their outputs are translated into concrete resource management actions, and how their failures are detected. This integration challenge is one of the main reasons why AI-based scheduling remains difficult to operationalize beyond controlled research environments.

4. Structural Trade-Offs in AI-Augmented Allocation

The introduction of AI into cloud scheduling creates a series of structural trade-offs that are not visible when algorithms are evaluated only on benchmark workloads. One fundamental trade-off is centralization versus decentralization. Centralized AI schedulers can exploit a global view of cluster state and workload demands, but they may become a bottleneck and a single point of failure. Decentralized AI agents can make faster local decisions, but they may lose coordination and produce globally suboptimal placements. Large-scale batch processing frameworks such as MapReduce demonstrated the value of centralized coordination for data-parallel workloads, while also revealing the overhead of global scheduling under highly variable task durations [16]. Later frameworks such as Spark introduced more flexible in-memory abstractions and improved scheduling for iterative workloads, but retained centralized coordination for many decisions [17].

A second trade-off involves model complexity and operational interpretability. Highly expressive models can capture subtle workload interactions, but they may be difficult for operators to understand, validate, or debug. In production environments, operators often need to explain why a workload was migrated or why a service experienced degraded performance. If the scheduling decision comes from a black-box model, diagnosing failures becomes harder. This can reduce trust and increase the time to recover from incidents. One practical response is to embed AI-based recommendations within a decision pipeline that retains explainable fallback policies. Another response is to use interpretable model surrogates or to expose confidence estimates alongside scheduling suggestions. These mechanisms do not eliminate the trade-off, but they can make learned policies more accountable in high-stakes operational settings.

A third trade-off is responsiveness versus stability. Cloud workloads can change rapidly, and schedulers that respond too aggressively may cause oscillations, unnecessary migrations, and resource thrashing. Conversely, schedulers that respond too slowly may miss opportunities to consolidate workloads or avoid congestion. AI-based schedulers are not immune to this challenge. A learned policy may overfit to transient patterns and issue frequent reallocations

that increase overhead. Stability can be improved by adding hysteresis, smoothing state representations, or applying penalties for unnecessary changes. However, these interventions can also reduce the efficiency gains that motivated AI adoption. The appropriate balance depends on workload characteristics, service-level objectives, and the cost of disruptive actions such as live migration.

Energy-proportional computing further complicates the trade-off landscape because the relationship between utilization and power draw is not linear across all hardware components [18]. AI-based allocation can help by predicting which hosts can be powered down or placed in low-power states without violating performance contracts. However, aggressive energy optimization may conflict with availability and responsiveness if it reduces spare capacity too much. System designers must therefore decide how to weight energy efficiency against risk in the reward structure of learning algorithms. These structural trade-offs indicate that AI-based resource allocation is not a purely technical optimization problem. It requires negotiated priorities among performance, cost, reliability, and sustainability.

5. Governance, Fairness, and Policy Implications

The deployment of AI-based schedulers in cloud infrastructures raises important governance and policy questions. In large-scale production environments, cluster management systems coordinate resources across many services, teams, and priorities. Systems such as Google's cluster manager have demonstrated that centralized control can achieve high utilization while supporting diverse workloads, but they also concentrate decision authority in a complex software layer [19]. When AI-based policies are added to such systems, accountability becomes more distributed. It may be unclear whether a service degradation was caused by an application change, a monitoring error, or a learned scheduling action. Governance mechanisms must therefore define who is responsible for reviewing AI decisions, how incidents are documented, and how operators can intervene when a learned policy behaves unexpectedly.

Fairness is another concern that extends beyond average performance metrics. In multitenant clouds, resource allocation decisions create winners and losers. A scheduling policy that optimizes global throughput may systematically delay workloads from lower-priority teams or constrain resources for latency-sensitive services. AI-based schedulers can inherit or amplify existing biases if they are trained on historical data that reflects past allocation inequalities. For example, if certain jobs were frequently preempted or throttled in the training data, the learned policy may continue to disadvantage similar jobs. Fairness-aware scheduling requires explicit attention to how resources are distributed across tenants, service classes, and geographic regions. It also requires evaluation metrics that capture tail behavior, not only average efficiency.

Policy implications include data governance, model provenance, and regulatory compliance. AI-based schedulers rely on telemetry that may contain sensitive information about workloads, users, or data flows. Organizations must ensure that data collection, storage, and use comply with privacy and security requirements. Model provenance is important for auditing decisions, reproducing failures, and demonstrating due diligence. If a learned scheduler causes a severe service disruption, operators may need to reconstruct the state that led to the decision. This requires logging not only the action but also the model version, input features, confidence levels, and safety constraints. Policy frameworks should also address the conditions under which autonomous decisions are acceptable. In some settings, human approval may be required for large-scale migrations, capacity reductions, or changes to critical services. In

other settings, autonomy may be necessary to meet latency requirements. Defining these boundaries is a sociotechnical task that involves infrastructure engineers, legal teams, and business stakeholders.

6. Infrastructure, Deployment, and Robustness Issues

Deploying AI-based scheduling algorithms in production cloud infrastructure introduces a distinct set of engineering challenges. The first challenge is data quality. Learned policies depend on consistent, accurate, and timely telemetry from hosts, containers, network devices, and application monitors. Missing data, clock skew, sensor noise, and heterogeneous labeling practices can degrade model performance in ways that are difficult to detect. Production systems therefore need robust data pipelines that validate, normalize, and timestamp observations before they reach the scheduler. The second challenge is serving latency. Some scheduling decisions must be made in milliseconds, while others can tolerate batch evaluation. This affects model architecture, hardware acceleration, and the placement of inference services. The third challenge is safe rollout. AI-based schedulers should be introduced through shadow evaluation, canary deployments, or constrained action spaces so that unexpected behavior can be observed before it affects critical workloads.

Autoscaling systems in large-scale production have shown that workload-aware automation can improve resource efficiency while reducing manual operations [20]. However, such systems also illustrate the importance of safeguards. A learned autoscaler may overprovision resources to avoid risk, eroding cost savings, or underprovision them to improve utilization, increasing the chance of service violations. Robustness requires monitoring for distribution shift, detecting when the relationship between workload features and performance changes, and retraining or fallback mechanisms. It also requires testing under adversarial conditions, including sudden demand spikes, partial infrastructure failures, and coordinated changes across multiple services. AI-based schedulers that perform well on average may still fail under tail events, and tail events are often the most consequential for user experience and service reliability.

Another deployment issue is the coupling between AI scheduling and existing cluster management abstractions. Cloud schedulers typically operate within hierarchical or multischeduler architectures. An AI component may be responsible only for container placement, while a separate autoscaler adjusts replica counts, and a network controller manages bandwidth. These components interact in ways that are difficult to model. A change in one layer can alter the state observed by another, leading to feedback loops or oscillatory behavior. Robust AI-based resource allocation must therefore be designed with explicit coordination boundaries and shared state models. It should also include run-time verification mechanisms that compare learned decisions against safe operating envelopes. Such mechanisms can prevent a learned policy from taking actions that violate hard constraints, even if those constraints were underrepresented during training.

7. Sustainability and Operational Perspectives

Sustainability has become an important objective for cloud resource allocation. Data centers consume substantial amounts of electricity, and their carbon footprint depends on both total energy use and the carbon intensity of the electricity supply. AI-based schedulers can contribute to sustainability by improving utilization, consolidating workloads onto fewer hosts, and shifting workloads to times or locations with lower carbon intensity. Energy-aware consolidation research has shown that dynamic migration can significantly reduce power

consumption [9]. However, sustainability cannot be reduced to a simple energy minimization objective. Aggressive consolidation may increase thermal stress, reduce hardware lifetime, or degrade performance in ways that trigger more resource consumption elsewhere. AI-based allocation systems should therefore incorporate sustainability as one dimension of a broader operational framework, not as an isolated optimization target.

Operational perspectives also include the human dimensions of automation. System administrators, capacity planners, and site reliability engineers must understand, trust, and maintain AI-based schedulers. If the system reduces operational workload but increases cognitive complexity, its overall value may be limited. Training and documentation become important, as do interfaces that allow operators to inspect pending decisions, query the rationale for a migration, and apply temporary overrides. Sustainability also has organizational dimensions. Long-lived cloud infrastructure requires policies that balance immediate cost savings against future capacity needs, hardware refresh cycles, and supply chain constraints. AI-based schedulers can support these decisions by modeling longer horizons, but only if their objectives are aligned with organizational priorities and their assumptions are transparent. This alignment is easier to achieve when AI systems are developed in close collaboration with the operations teams that will use them.

8. Future Directions and Research Agenda

Several future directions are likely to shape the evolution of AI-based cloud resource allocation. First, robust learning under distribution shift remains an open problem. Cloud workloads change as applications evolve, user behavior shifts, and hardware configurations are updated. AI schedulers must detect these shifts and update their policies without causing instability. Second, cross-layer coordination between compute, storage, networking, and power management deserves deeper study. Current AI-based schedulers often optimize one layer while treating others as fixed. Integrated learning across layers could yield better system-level outcomes but requires new modeling approaches and coordination protocols. Third, explainability and verification need to advance beyond post hoc feature attribution. Operators require formal or semi-formal guarantees about how a learned scheduler will behave under known constraints. Research on safe reinforcement learning and constrained policy optimization may help, but translating those methods to production-scale systems remains difficult.

Fourth, fairness and governance require more systematic treatment. Researchers should develop metrics that capture distributional effects across tenants and service classes, and should evaluate AI schedulers not only on aggregate efficiency but also on equity and procedural fairness. Fifth, federated or multi-agent learning may be needed for decentralized cloud environments that span multiple data centers, edge sites, and administrative domains. Such approaches must address privacy, communication overhead, and incentive alignment. Sixth, sustainability-aware scheduling should include carbon-aware policies that respond to real-time grid signals, workload flexibility, and energy storage constraints. Finally, longitudinal field studies of AI-based schedulers in production would provide valuable evidence about their operational impact. Academic research has produced many promising algorithmic prototypes, but fewer studies examine how these systems behave over months or years under real operational conditions. Closing this gap will require closer collaboration between academia, cloud providers, and enterprise infrastructure teams.

9. Conclusion

AI-based scheduling algorithms have the potential to improve cloud resource allocation by learning from operational telemetry, adapting to workload changes, and balancing multiple objectives that are difficult to encode in static heuristics. However, their successful deployment is not simply a matter of predictive accuracy or algorithmic sophistication. System-level factors such as architectural placement, governance structures, fairness requirements, robustness under uncertainty, deployment safeguards, and sustainability objectives determine whether learned schedulers produce meaningful and lasting benefits. The analysis presented in this paper highlights the structural trade-offs between centralized and decentralized control, responsiveness and stability, model expressiveness and interpretability, and efficiency and equity. It also shows that AI-based allocation must be embedded within a broader sociotechnical system that includes monitoring, accountability, human oversight, and policy compliance. Future research should move beyond isolated algorithm benchmarks and toward integrated evaluations that capture the full operational life cycle of AI-based resource management. By doing so, the cloud computing community can develop schedulers that are not only intelligent but also trustworthy, sustainable, and aligned with the long-term needs of the organizations and societies that depend on cloud infrastructure.

References

1. Mell, P., & Grance, T. (2011). The NIST definition of cloud computing (NIST Special Publication 800-145). National Institute of Standards and Technology.
<https://doi.org/10.6028/NIST.SP.800-145>
2. Buyya, R., Yeo, C. S., Venugopal, S., Broberg, J., & Brandic, I. (2009). Cloud computing and emerging IT platforms: Vision, hype, and reality for delivering computing as the 5th utility. *Future Generation Computer Systems*, 25(6), 599–616.
<https://doi.org/10.1016/j.future.2008.12.001>
3. Armbrust, M., Fox, A., Griffith, R., Joseph, A. D., Katz, R., Konwinski, A., Lee, G., Patterson, D., Rabkin, A., Stoica, I., & Zaharia, M. (2010). A view of cloud computing. *Communications of the ACM*, 53(4), 50–58. <https://doi.org/10.1145/1721654.1721672>
4. Jennings, B., & Stadler, R. (2015). Resource management in clouds: Survey and research challenges. *Journal of Network and Systems Management*, 23(3), 567–619.
<https://doi.org/10.1007/s10922-014-9307-7>
5. Sutton, R. S., & Barto, A. G. (2018). *Reinforcement learning: An introduction* (2nd ed.). MIT Press.
6. Mao, H., Alizadeh, M., Menache, I., & Kandula, S. (2016). Resource management with deep reinforcement learning. *Proceedings of the 15th ACM Workshop on Hot Topics in Networks*, 50–56. <https://doi.org/10.1145/3005745.3005750>
7. Zhang, Q., Cheng, L., & Boutaba, R. (2010). Cloud computing: State-of-the-art and research challenges. *Journal of Internet Services and Applications*, 1(1), 7–18.
<https://doi.org/10.1007/s13174-010-0007-6>
8. Calheiros, R. N., Ranjan, R., Beloglazov, A., De Rose, C. A. F., & Buyya, R. (2011). CloudSim: A toolkit for modeling and simulation of cloud computing environments and evaluation of resource provisioning algorithms. *Software: Practice and Experience*, 41(1), 23–50. <https://doi.org/10.1002/spe.995>

9. Beloglazov, A., & Buyya, R. (2012). Optimal online deterministic algorithms and adaptive heuristics for energy and performance efficient dynamic consolidation of virtual machines in cloud data centers. *Concurrency and Computation: Practice and Experience*, 24(13), 1397–1420. <https://doi.org/10.1002/cpe.1867>
10. Xiao, Z., Song, W., & Chen, Q. (2013). Dynamic resource allocation using virtual machines for cloud computing environment. *IEEE Transactions on Parallel and Distributed Systems*, 24(6), 1107–1117. <https://doi.org/10.1109/TPDS.2012.273>
11. Chen, T., & Guestrin, C. (2016). XGBoost: A scalable tree boosting system. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 785–794. <https://doi.org/10.1145/2939672.2939785>
12. LeCun, Y., Bengio, Y., & Hinton, G. (2015). Deep learning. *Nature*, 521(7553), 436–444. <https://doi.org/10.1038/nature14539>
13. Silver, D., Huang, A., Maddison, C. J., Guez, A., Sifre, L., van den Driessche, G., Schrittwieser, J., Antonoglou, I., Panneershelvam, V., Lanctot, M., Dieleman, S., Grewe, D., Nham, J., Kalchbrenner, N., Sutskever, I., Lillicrap, T., Leach, M., Kavukcuoglu, K., Graepel, T., & Hassabis, D. (2016). Mastering the game of Go with deep neural networks and tree search. *Nature*, 529(7587), 484–489. <https://doi.org/10.1038/nature16961>
14. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
15. Mirhoseini, A., Pham, H., Le, Q. V., Steiner, B., Larsen, R., Zhou, Y., Kumar, N., Norouzi, M., Bengio, S., & Dean, J. (2017). Device placement optimization with reinforcement learning. *Proceedings of the 34th International Conference on Machine Learning*, 2430–2439.
16. Dean, J., & Ghemawat, S. (2008). MapReduce: Simplified data processing on large clusters. *Communications of the ACM*, 51(1), 107–113. <https://doi.org/10.1145/1327452.1327492>
17. Zaharia, M., Chowdhury, M., Franklin, M. J., Shenker, S., & Stoica, I. (2010). Spark: Cluster computing with working sets. *Proceedings of the 2nd USENIX Conference on Hot Topics in Cloud Computing*, 10.
18. Barroso, L. A., & Hölzle, U. (2007). The case for energy-proportional computing. *Computer*, 40(12), 33–37. <https://doi.org/10.1109/MC.2007.443>
19. Verma, A., Pedrosa, L., Korupolu, M., Oppenheimer, D., Tune, E., & Wilkes, J. (2015). Large-scale cluster management at Google with Borg. *Proceedings of the Tenth European Conference on Computer Systems*, 1–17. <https://doi.org/10.1145/2741948.2741964>
20. Rzadca, K., Findeisen, P., Swiderski, J., Zych, P., Broniek, P., Kusmieriek, J., Nowak, P., Strack, B., Witusowski, P., Hand, S., & Wilkes, J. (2020). Autopilot: Workload autoscaling at Google. *Proceedings of the Fifteenth European Conference on Computer Systems*, 1–16. <https://doi.org/10.1145/3342195.3387524>