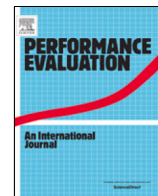




Contents lists available at ScienceDirect

Performance Evaluation

journal homepage: www.elsevier.com/locate/peva

The case for phase-aware scheduling of parallelizable jobs

Benjamin Berg^{*}, Justin Whitehouse, Benjamin Moseley, Weina Wang,
Mor Harchol-Balter

Carnegie Mellon University, United States of America

ARTICLE INFO

Article history:
Available online xxxx

Keywords:
Performance modeling
Parallel scheduling
Server allocation
Databases

ABSTRACT

Parallelizable jobs typically consist of multiple phases of computation, where the job is more parallelizable in some phases and less parallelizable in others. For example, in a database, a query may consist of a highly parallelizable table scan, followed by a less parallelizable table join. In the past, this phase-varying parallelizability was summarized by a single sub-linear speedup curve that measures a job's average parallelizability over its entire lifetime. Today, however, modern systems have fine-grained knowledge of the exact phase each job is in at every moment in time. Unfortunately, these systems do not fully leverage this real-time feedback when scheduling parallelizable jobs. Theory has failed to produce practical phase-aware scheduling policies, and thus scheduling in current systems is largely heuristic.

A phase-aware scheduling policy must decide, given its knowledge of each job's current phase, how many servers or cores to allocate to each job in the system at every moment in time. This paper provides the first stochastic model of a system processing parallelizable jobs composed of phases. Using our model, we derive an optimal phase-aware scheduling policy that minimizes the mean response time across jobs. Our provably optimal policy, Inelastic-First (IF), gives strict priority to jobs that are currently in less parallelizable phases. We validate our theoretical results using a simulation of a database running queries from the Star Schema Benchmark. We compare IF to a range of policies from both systems and theory, and show that IF reduces mean response time by up to a factor of 3.

© 2021 Elsevier B.V. All rights reserved.

1. Introduction

Parallelizable workloads are ubiquitous and appear across a diverse array of modern computer systems. Data centers, supercomputers, machine learning clusters, distributed computing frameworks, and databases all process jobs designed to be parallelized across many servers or cores. Unlike the jobs in more classical models, such as the M/G/k queue, that each run on a single server, parallelizable jobs are capable of running on multiple servers simultaneously. When a job is parallelized across additional servers or cores, the job receives a speedup and can be completed more quickly.

When scheduling parallelizable jobs, a *scheduling policy* must decide *how to best allocate servers or cores among the jobs in the system at every moment in time*. This paper describes and analyzes scheduling policies for systems that process an online *stream* of incoming parallelizable jobs. Given a set of K servers, we will derive scheduling policies that minimize the *mean response time* across jobs — the average time from when a job arrives to the system until it is completed.

^{*} Corresponding author.

E-mail addresses: bsberg@cs.cmu.edu (B. Berg), jwhiteho@andrew.cmu.edu (J. Whitehouse), moseleyb@andrew.cmu.edu (B. Moseley), weinaw@cs.cmu.edu (W. Wang), harchol@cs.cmu.edu (M. Harchol-Balter).

<https://doi.org/10.1016/j.peva.2021.102246>

0166-5316/© 2021 Elsevier B.V. All rights reserved.

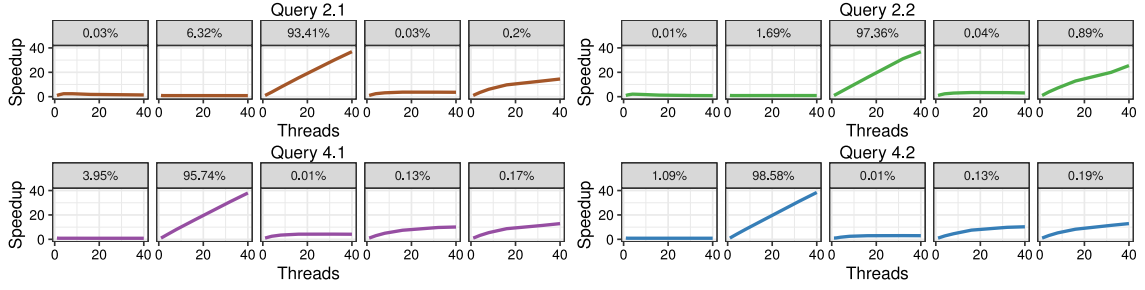


Fig. 1. Speedup functions for each phase of four queries from the Star Schema Benchmark. Queries were executed using the NoisePage database [1]. Phases are either elastic (highly parallelizable) or inelastic (highly sequential). The percentages denote the fraction of time spent in each phase when the query was run on a single core. Despite the queries spending most of their time in elastic phases, the overall speedup function of each query is highly sublinear due to Amdahl's law.

The difficulty in scheduling parallelizable jobs arises largely from the fact that a job's parallelizability is not constant over time. Across a wide variety of systems, jobs typically consist of multiple *phases*, each of which has its own scalability characteristics.

For example, in databases, a single query often alternates between highly parallelizable phases and non-parallelizable phases. Specifically, modern databases translate queries into a pipeline composed of multiple phases corresponding to different database operations [1]. A phase that corresponds to a sequential table scan is *elastic*, capable of perfectly parallelizing and completing k times faster when run on k cores. On the other hand, a phase corresponding to a table join is *inelastic*, receiving a severely limited speedup from additional cores. Fig. 1 shows that this phenomenon holds for a variety of queries from the Star Schema Benchmark [2].

Our problem

In practice, many system schedulers are aware of each job's current phase [3,4]. Databases explicitly invoke the elastic and inelastic phases of a database query during query execution [5]. Cluster schedulers [6], distributed computing platforms (e.g. Hadoop [7] and Apache Spark [8]), distributed machine learning frameworks [9], and supercomputers all process jobs composed of a mixture of highly parallelizable and highly sequential phases.

Although the above systems can track the current phase of each running job, they do not effectively leverage this information to make optimal scheduling decisions. In this paper, we address the problem of *phase-aware scheduling* — using the available phase information to allocate resources efficiently across jobs. Given a stream of parallelizable jobs composed of multiple phases, our goal is to design scheduling policies which decide, at every moment in time, how many cores or servers to allocate to each job.

Why phase-aware scheduling has not been solved

The systems community, theoretical computer science (TCS) community, and stochastic performance modeling community have all done significant work on the problem of parallel job scheduling. However, we will show that phase-aware scheduling policies can improve both existing theoretical results and state-of-the-art systems schedulers.

One approach the systems community uses is to defer scheduling decisions to the user by relying on reservation-based systems [10–12]. Here, users reserve the number of cores or servers on which they want to run their jobs. Unfortunately, studies show that users tend to conservatively over-provision resources, leading to suboptimal resource allocations [6,11].

Although phase-aware systems schedulers do exist, they often make suboptimal scheduling decisions. For instance, database schedulers use phase information to avoid over-allocating cores to queries which are in an inelastic phase, but otherwise process queries in first-come-first-served (FCFS) order [13]. We refer to this policy as *Phase-Aware FCFS* (PA-FCFS) to distinguish it from a naive FCFS policy that over-allocates to jobs in inelastic phases. We will show that PA-FCFS can be far from optimal.

The approach of the TCS community has been to analyze the problem through the lens of competitive analysis, where it is assumed that the arrival sequence of jobs is chosen adversarially. This work either assumes that jobs consist of phases with different degrees of scalability or that each job is encoded as a directed acyclic graph (DAG) [14,15]. In these adversarial settings, the result from [16] provides strong lower bounds on the achievable competitive ratio. In particular, no scheduling policy can perform within a constant factor of the optimal policy in the worst case. The TCS community has also found policies that match these lower bounds, such as the EQUI policy [15,17] that divides servers evenly between all jobs currently in the system. This has led the TCS community to conclude that the problem of scheduling parallelizable jobs is solved, even though these policies frequently perform worse than PA-FCFS (see Section 7).

The stochastic community has thus far largely assumed that all jobs follow the same, single speedup function that dictates how parallelizable the jobs are [18]. However, this work has not addressed the scheduling of jobs whose parallelizability changes over time.

Optimal phase-aware scheduling

In summary, although real-world systems process jobs composed of phases, and these systems are often aware of the current phase of each job, phase-aware scheduling remains an open problem. Hence, our first contribution is a stochastic model of jobs composed of multiple phases with different levels of parallelizability. Under this model, we derive a provably optimal scheduling policy. This optimal policy, IF, is non-obvious and greatly outperforms both the PA-FCFS policy used in real systems and the EQUI policy proposed in the worst-case literature. Because our model makes some simplifying assumptions, we validate IF's performance through a range of simulations, including a simulation of a database running queries from the Star Schema Benchmark [2].

Contributions of this paper

- We first present a novel model of parallelizable jobs composed of elastic and inelastic phases where the scheduler knows, at all times, what phase a job is in. Our model is far more general than prior work from the stochastic community which has assumed that all jobs follow the same, single speedup function.
- We prove that the *Inelastic First* (IF) policy, which *defers parallelizable work* by giving strict priority to jobs which are in an inelastic phase, is optimal under our model. Because the proof of optimality requires a complex coupling argument, we break this claim down by considering special cases which are easier to understand. We begin by proving the optimality of IF in simpler models in Section 5 before proving our more general claim in Section 6.
- In Section 7, we perform an extensive simulation-based performance evaluation to illustrate that IF outperforms a range of scheduling policies. Even in settings that violate the assumptions of our model, IF can perform nearly 30% better than the PA-FCFS policy used in modern databases and a factor of 3 better than the EQUI policy advocated by the TCS community.
- Lastly, in Section 8, we perform a case study on scheduling in databases where queries consist of elastic and inelastic phases. In this setting, the scheduler sometimes has additional information about each query beyond just the query's current phase. We show how to generalize IF to leverage this additional information and improve upon state-of-the-art database scheduling by roughly 50% in simulation.

2. Prior work

It is easiest to categorize the prior theoretical work on scheduling parallelizable jobs by the model of parallelism considered. We will therefore discuss the different theoretical models of parallelism before considering prior work from the systems community on scheduling parallelizable jobs.

Jobs with parallelizable phases

The closest theoretical work to ours comes from the worst-case scheduling community [15,19–21]. This work similarly considers the problem of scheduling parallelizable jobs composed of phases of differing parallelizability. Due to the worst-case nature of the analysis, this work is forced to either consider an offline problem where all jobs arrive at time 0 [19], or to rely on resource augmentation¹ [15,20,21] to provide an algorithm which is within a (potentially large) constant factor of the optimal policy. This work concludes that the EQUI policy, as well as a generalization of it, is constant competitive given a small constant resource augmentation.

A related work from the parallel scheduling community recognizes that jobs have elastic and inelastic phases [22]. However, for analytical tractability, [22] assumes that jobs consist of only a *single* phase, and are therefore either fully elastic or fully inelastic. Even in this limited setting, [22] requires that the inelastic jobs are smaller on average than the elastic jobs. By contrast, our model allows each job to have any number of phases, with different jobs having different numbers of phases. Furthermore, our model does not make any assumptions about the relative sizes of elastic and inelastic phases.

Jobs with speedup curves

Other theoretical work has also considered a model where, instead of consisting of phases, each job follows a single *speedup function*, $s(k)$, that describes the speedup a job receives from running on k servers. Here, $s(k)$ is some positive, concave, non-decreasing function. Work using this model from the worst-case scheduling literature finds that when job sizes are known, a generalization of EQUI is $O(\log p)$ -competitive with the optimal policy, where p is the ratio of the largest job size to the smallest job size [17]. Moreover, EQUI is again shown to be constant competitive with constant resource augmentation [20,21]. In an analogous result using this model from the performance modeling community, [18] finds that EQUI is the optimal policy when job sizes are unknown and exponentially distributed.

¹ Resource augmentation analysis is a relaxation of competitive analysis that, for some $s > 1$, compares an algorithm using speed s processors against the optimal policy using speed 1 processors.

Overall, the general consensus from both the worst-case scheduling community and the performance modeling community is that EQUI should be used to achieve good or possibly optimal mean response time. However, as we will see, EQUI is far from optimal when jobs are composed of elastic and inelastic phases (see Fig. 7). This discrepancy is due to the overly pessimistic nature of the prior theory work, which assumes that the system cannot determine how parallelizable a job is at each moment in time. We assume that the scheduler knows whether a job is in an elastic phase or an inelastic phase, which is reasonable for a wide range of systems [23–26]. As a result, we are able to provide an optimal phase-aware scheduling policy that consistently outperforms EQUI.

DAG jobs

A separate branch of theoretical work that developed concurrently with the above models considers every job as consisting of a set of tasks with precedence constraints specified by a Directed Acyclic Graph (DAG). In this model, introduced in [27], a task can only run on a single server, but any two tasks that do not share a precedence relationship can be run in parallel. Much of the work in this area is concerned with how to efficiently schedule a *single* DAG job onto a set of servers [27–29]. When multiple DAG jobs arrive over time, there are strong lower bounds on the competitive ratio of any online algorithm for mean response time [16]. Recently, [14] considered the online problem of scheduling a stream of DAG jobs to minimize the worst case mean response time. Using a resource augmentation argument, they show that EQUI and its generalization are constant competitive when given constant resource augmentation.

Systems literature

The need to schedule jobs with sublinear speedup functions has been corroborated across a wide range of systems. Perhaps most famously, the computer architecture community identified Amdahl's law [30] around the advent of multicore architectures. The problem of scheduling parallelizable jobs is similarly known in the context of data center scheduling [6], supercomputing [31,32], distributed machine learning [9], databases [33], and distributed computing frameworks such as MapReduce [24,34]. Existing schedulers in these contexts are highly dependent on heuristics [6,23,26,35], often require significant parameter tuning, and do not provide formal guarantees about performance. Our goal is to improve upon these state-of-the-art heuristic policies by providing practical policies with provably optimal or near-optimal performance.

3. Model

In this section, we present a model of jobs composed of distinct phases running in a system consisting of K homogeneous servers.

Multi-phase jobs

We begin by noting that in a wide range of systems applications, job phases are either highly parallelizable or highly sequential. Fig. 1 shows that database queries often follow this pattern. A similar phenomenon applies in systems using a map-reduce paradigm [24] where parallelizable map stages are interlaced with sequential reduce stages. Machine learning training jobs also consist of highly parallelizable iterations of distributed gradient descent followed by a sequential step that coalesces the results on a central parameter server [9]. Hence, while job phases could potentially experience intermediate parallelizability, we will consider the highly practical case where job phases are either *elastic*, perfectly parallelizable, or *inelastic*, totally sequential.

To model the duration of each job phase, we define a phase's *inherent size* to be the amount of time it takes the phase to complete when run on a single server. For analytical tractability, we will assume that inelastic phase sizes are distributed as $\text{Exp}(\mu_I)$ and elastic phase sizes are distributed as $\text{Exp}(\mu_E)$, and that all phase sizes are independently distributed. Although the scheduler often knows the current phase of each job in the system, it is less common in real systems for the scheduler to know the full sequence of phases comprising each job, or the size of each phase. Hence, we will generally assume that the scheduler knows the current phase of each job, but that the scheduler does not know the future phases or any of the phase sizes of a job. In Section 8, we consider scheduling in databases, where it is common for the database to have additional information about each job's phases and phase sizes.

Only elastic phases can be parallelized across multiple servers. An elastic phase of size x , when run on k servers, takes $\frac{x}{k}$ time to complete. Equivalently, we model the running time of an elastic phase on k servers as a random variable which is distributed as $\text{Exp}(k\mu_E)$. By contrast, an inelastic phase cannot be parallelized and runs on at most one server at any moment in time.

Because the sizes of a job's phases are assumed to be exponentially distributed and unknown to the system, we model a multi-phase job via a continuous-time Markov chain, as shown in Fig. 2. We model each job via a Markov chain consisting of three states: an E state that denotes that the job is in an elastic phase, an I state that denotes that the job is in an inelastic phase, and an absorbing state, C , that denotes that the job has been completed. Each arriving job can either start in the E state or in the I state. We assume that a job can only transition to the completion state from the inelastic state. This is realistic for a wide range of systems where the results of a parallel computation must be sequentially coalesced



Specifically, if there are i jobs in the system that are in an inelastic phase, then IF allocates $\min\{i, K\}$ servers to these jobs. IF then allocates any remaining servers to a job in an elastic phase if such a job exists, otherwise these extra servers remain idle. Deferring parallelizable work can increase system efficiency by keeping flexible elastic phases in the system, ensuring that a policy can utilize all K servers.

We show that IF is optimal with respect to minimizing mean response time. Intuitively, IF is a good policy because it is able to both favor jobs with smaller expected remaining sizes, and defer parallelizable work. Observe that IF does not require any knowledge of the job parameters (μ_I , μ_E , and q). Thus, optimally scheduling multi-phase jobs can be done regardless of whether these parameters are known to the system.

Arrival processes and metrics

We allow for an arbitrary arrival process. To be precise, we first define an *arrival time sequence* as two fixed, infinite sequences, $(t_n)_{n \geq 1}$ and $(\ell_n)_{n \geq 1}$, where t_n is the time at which the n th job arrives and $\ell_n \in \{E, I\}$ denotes whether the arriving job begins with either an E phase or an I phase. We define an *arrival time process* as a distribution over arrival sequences.

We define the *response time* of the n th job under policy π to be the time from when the job arrives until it completes. We denote this quantity by the random variable $T_\pi^{(n)}$. We let T_π denote the steady-state response time whenever this quantity exists.

As an example, consider the case where the arrival time process is a Poisson process with rate λ and each job starts with an E phase with probability r_E and with an I phase with probability r_I . Then we can define the system load as:

$$\rho = \text{System load} = \frac{\lambda \cdot \mathbb{E}[\text{Job size}]}{K},$$

where

$$\mathbb{E}[\text{Job size}] = r_E \cdot \mathbb{E}[\text{Job size if start in } E] + r_I \cdot \mathbb{E}[\text{Job size if start in } I].$$

In this setting, if $\rho < 1$, the steady-state mean response time under policy π exists and is denoted by $\mathbb{E}[T_\pi]$.

Stochastically minimizing the number of jobs in system

We show that IF minimizes the steady-state mean response time across jobs. To show this, we prove a series of claims about the number of jobs in the system at any point in time. Namely, we argue that IF stochastically maximizes the number of jobs completed by any point in time. This is equivalent to saying IF stochastically minimizes the number of jobs in system at any point in time.

To reason about the number of completions by time t , we will count the number of elastic and inelastic transitions as well as the number of job completions. We define $C_\pi(t)$ to be the number of job completions by time t under policy π . We define $I_\pi(t)$ (and $E_\pi(t)$) to be the number of inelastic (resp. elastic) transitions under policy π by time t . Finally, we define $I_\pi(s, t)$ to be the number of inelastic transitions under π on the interval $(s, t]$ and we define $E_\pi(s, t)$ and $C_\pi(s, t)$ analogously.

With respect to the number of jobs in system, let $N_\pi(t)$ denote the number of jobs present at time t , under policy π . We define $N_\pi^E(t)$ to be the number of jobs in an elastic phase at time t under π and we define $N_\pi^I(t)$ to be the number of jobs in an inelastic phase at time t under π .

4. Overview of theorems

In this section, we provide an overview of the theoretical results in Sections 5 and 6.

4.1. Main result

We first state the main theorem in full generality. At a high level, the theorem states that IF is the most effective policy in terms of completing jobs. More specifically, we show that the number of jobs completed by any point in time under IF stochastically dominates the number of jobs completed by the same time under any other algorithm.

Theorem 1. Consider a K server system serving multi-phase jobs. The policy IF stochastically maximizes the number of jobs completed by any point in time. Specifically, for a policy A , let $C_A(t)$ denote the number of jobs completed by time t and let $N_A(t)$ denote the number of jobs in the system at t . Then under any arbitrary arrival time process, $C_{\text{IF}}(t) \geq_{st} C_A(t)$ for all times $t \geq 0$. Consequently, $N_{\text{IF}}(t) \leq_{st} N_A(t)$ for all times $t \geq 0$.

We can leverage Theorem 1 to derive far-reaching results about job response time. In particular, if the arrival time process is a renewal process,² we can show that IF minimizes the steady-state mean response time. We formalize this idea in the following immediate corollary of Theorem 1.

² By a renewal process, we mean the inter-arrival times $t_n - t_{n-1}$ are i.i.d., and that the initial phases of jobs p_n are i.i.d. as well.

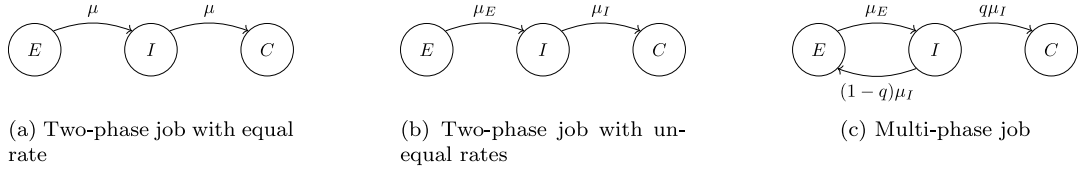


Fig. 4. The three job structures we consider. E refers to the elastic state, I refers to the inelastic state, and C refers to the completion state. In Fig. 4(a), jobs just have two phases, both with inherent size distributed as $\text{Exp}(\mu)$. In Fig. 4(b), jobs still have two phases. Phase I has size distributed as $\text{Exp}(\mu_I)$, and phase E has size distributed as $\text{Exp}(\mu_E)$. In Fig. 4(c), we add in potential transitions from an inelastic phase to an elastic phase.

Corollary 2. Suppose the same system setup as in Theorem 1. For any arbitrary policy A , let T_A be the steady-state job response time when it exists. If the arrival time process is a renewal process, then $\mathbb{E}[T_{\text{IF}}] \leq \mathbb{E}[T_A]$.

Proof. By Theorem 1, we know IF stochastically minimizes the number of jobs in the system at any point in time. Since the arrival time process is a renewal process, this implies that IF minimizes the steady-state mean number of jobs in the system. By Little’s law, minimizing the mean number of jobs in the system suffices for minimizing the steady-state mean response time. $\square$

Theorem 1 and its corollary show that IF succeeds by both *deferring parallelizable work* and working on jobs with smaller expected remaining sizes. Specifically, while elastic phases can be completed more quickly by parallelizing across all servers, there are benefits to keeping elastic phases in the system. These elastic phases are flexible and can ensure that all K servers remain utilized. It is also possible to allocate some servers to inelastic phases without significantly increasing the runtime of an elastic phase. Furthermore, jobs in inelastic phases have smaller expected remaining sizes. Hence, deferring parallelizable work also results in favoring shorter jobs. For these reasons, the optimal policy, IF, defers as much parallelizable work as possible without over-allocating to inelastic phases.

Remark 3. One might assume that IF benefits not from deferring parallelizable work, but rather from how we have defined our model, where jobs in inelastic phases have smaller expected remaining sizes. However, as we show in Section 8, simply favoring jobs with smaller remaining sizes (as done by the PA-SRPT policy) is not nearly as important as deferring parallelizable work in real-world settings.

4.2. How We Prove Theorem 1

We now provide a road map for how we prove Theorem 1. At a high level, we note that it suffices to find a coupling between two systems, one running IF and one running an arbitrary policy A , under which $C_{\text{IF}}(t) \geq C_A(t)$, $\forall t \geq 0$. However, finding such a coupling is difficult due to the complicated job structure in Fig. 2. In particular, the inherent size distributions are different between the two phases and jobs are composed of an unknown number of elastic and inelastic phases.

We therefore begin by considering several simpler job structures, as seen in Fig. 4. The simplest job structure has elastic and inelastic phases with the same size distribution, and does not allow transitions from an inelastic phase to an elastic phase. We then add the complexities gradually back to the model. In each case, we argue that studying the number of inelastic transitions suffices to understand the total number of job completions. Recalling that $I_A(t)$ is the number of inelastic transitions under A by time t , we prove that, for any policy A , there exists a coupling under which $I_{\text{IF}}(t) \geq I_A(t)$ for all $t \geq 0$. We then argue that, under this coupling, $C_{\text{IF}}(t) \geq C_A(t)$ for all $t \geq 0$.

In Section 5.1, we start with the simplest job structure as shown in Fig. 4(a). In this structure, jobs consist of a single elastic phase followed by a single inelastic phase. Moreover, we assume the inherent sizes of both the elastic and inelastic phases are identically distributed as $\text{Exp}(\mu)$. We refer to such jobs as *two-phase jobs with equal rates*. We are able to couple two systems experiencing jobs of this structure by (1) having them experience the same sequence of arrivals and (2) splitting time into roughly uniform chunks of length $\text{Exp}(K\mu)$. At the end of each chunk of time, the systems will both potentially experience a job transition. By splitting time into “busy” and “idle” periods under this coupling (as defined in the proof of Lemma 4), we prove the desired result.

We then consider the slightly more complicated job structure shown in Fig. 4(b) in Section 5.2. In this job structure, jobs again consist of a single elastic phase followed by a single inelastic phase. However, we drop the assumption that the inherent sizes of elastic and inelastic phases are identically distributed. We refer to such jobs as *two-phase jobs with unequal rates*. Although having unequal rates between phases complicates the splitting of time into roughly equal blocks, we work around this by leveraging a trick called *uniformization*. More specifically, we reformulate this more general job structure as a Markov chain in which the elastic and inelastic phases have the same inherent size distribution, but with some additional self-loop transitions added to the chain. We then expand our existing coupling argument by coupling the transition outcomes of the two systems.

Finally, we consider the general job structure as shown in Fig. 4(c) in Section 6. In this job structure, jobs have alternating elastic and inelastic phases, each with a different service rate. We refer to such jobs as *multi-phase jobs*. This case seems quite different from the previous settings, since now an inelastic transition can produce an elastic phase. However, we show that such a transition is equivalent to a job completion followed immediately by an arrival of a job beginning with an elastic phase. Using this argument, we show how a coupling in the general case follows from our coupling in the cases with two-phase jobs.

The above arguments all leverage the fact that, for a given sample path under our model, the policy that has completed more inelastic phases by time t will also have completed more jobs. This is a direct result of our assumption that the final phase of each job will be an inelastic phase. In the case where jobs may finish with an elastic phase, an optimal policy may give priority to a job in an elastic phase if the job has a sufficiently small expected remaining size. We examine the case where jobs finish with elastic phases in Section 7 and find that IF still frequently performs well.

5. Two-phase jobs

5.1. Two-phase jobs with equal rates

We first consider two-phase jobs with equal rates. These are jobs that consist of a single elastic phase followed by a single inelastic phase where both phases have inherent size distributed as $\text{Exp}(\mu)$, as illustrated in Fig. 4(a).

Lemma 4. Consider a K server system serving two-phase jobs with equal rates. Consider any policy A and let the policies IF and A start from the same initial conditions and have the same arrival time process. Then there exists a coupling between IF and A such that $I_{\text{IF}}(t) \geq I_A(t)$ and $C_{\text{IF}}(t) \geq C_A(t)$ for all $t \geq 0$, where $I_{\text{IF}}(t)$ (resp. $I_A(t)$) is the number of inelastic transitions by time t under IF (resp. A), and $C_{\text{IF}}(t)$ (resp. $C_A(t)$) is the number of jobs completed by time t under IF (resp. A).

The proof of Lemma 4 can be found in the Appendix. We describe the coupling used in the proof below in Section 5.1.1, as it serves as a building block for subsequent arguments.

Note that for two-phase jobs with equal rates, every inelastic job transition is also a completion, so we have $I_A(t) = C_A(t)$ for all times $t \geq 0$ under any policy A . Therefore, to prove Lemma 4, it suffices to construct a coupling under which $I_{\text{IF}}(t) \geq I_A(t)$ for all $t \geq 0$. Then the claim $C_{\text{IF}}(t) \geq C_A(t)$ follows directly.

5.1.1. Coupling IF and A

Let S_{IF} be the system running IF and S_A be the system running any arbitrary policy A . The high-level intuition of the coupling is as follows. Since both phases, inelastic and elastic, have inherent size $\text{Exp}(\mu)$, we can parse time into blocks of length $\text{Exp}(K\mu)$. At the end of each of these blocks, both systems will potentially experience a job transition. Outside of these points of time, no job transitions can occur. This simplifies the counting of job completions and inelastic transitions. Arrivals do not change the number of transitions or job completions, and hence we do not need assumptions on the arrival time process.

Job arrivals: We assume that the two systems, S_{IF} and S_A , have the same number of jobs in each phase at time 0 (for instance, seven jobs in an inelastic phase, and three jobs in an elastic phase). Formally, we assume that $N_{\text{IF}}^E(0) = N_A^E(0)$ and $N_{\text{IF}}^I(0) = N_A^I(0)$.

We fix an arrival time sequence that is shared between S_{IF} and S_A . Recall that an arrival sequence is a fixed sequence of arrival times $(t_n)_{n \geq 1}$ and a corresponding binary sequence $(\ell_n)_{n \geq 1}$, where t_n is the time the n th overall job arrival occurs in both systems and $\ell_n \in \{E, I\}$ determines which phase a job starts in.

Job transitions and departures: Suppose the current time is t . We generate a random variable $X \sim \text{Exp}(K\mu)$, that is shared by both systems. Suppose s is the next unrealized arrival time in the arrival sequence. If $s < t + X$, we allow the arrival to occur simultaneously into both systems. We then set $t \leftarrow s$, and return to the beginning of this paragraph. If $s > t + X$, then we set the current time to be $t \leftarrow t + X$, and then select one of the K servers uniformly at random³ (we select the same server in both systems). If a system is running a job in its inelastic phase on this randomly selected server, the job completes. Likewise, if the server is running a job in its elastic phase, the system experiences an elastic transition, producing an inelastic phase. Lastly, if the server selected is idling, nothing happens. This event (which may or may not result in a transition/departure) will be referred to as a *potential transition*. In general, a time where either an arrival or potential transition occurs will be referred to as an *event time*.

Additionally, if at time t the system is serving i inelastic jobs, we assume they are running on servers 1 through i . If an elastic job is being served, it is run on servers $i + 1$ through e , where e is some number less than or equal to K . The remaining servers are left idle.

³ For the sake of simplicity, we assume that jobs can only be allocated an integral number of servers. However, our result generalizes to the case where allocations are fractional. When allocations are fractional, we treat the servers as a continuous interval, $[0, K]$ and generate $U \sim \text{Unif}[0, K]$. The type of phase running at the corresponding point in the interval $[0, K]$ determines what type of transition occurs.

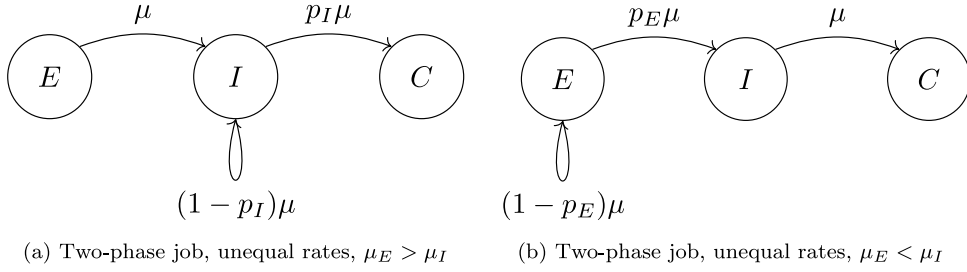


Fig. 5. Two cases of uniformizing two-phase jobs with unequal rates. In Fig. 5(a), $\mu_E > \mu_I$, so we take our dominating rate as $\mu := \mu_E$. We then take $p_I := \frac{\mu_I}{\mu}$ and set the total transition rate out of the inelastic state to be μ . With probability $1 - p_I$, after completing an inelastic phase, we immediately start another one. With probability p_I , the job completes and exits the system. Fig. 5(b) shows the analogous case where $\mu_E < \mu_I$.

5.2. Two-phase jobs with unequal rates

We now consider two-phase jobs with unequal rates, as illustrated in Fig. 4(b). In this case, the inherent sizes of elastic phases are distributed as $\text{Exp}(\mu_E)$, and the inherent sizes of inelastic phases are distributed as $\text{Exp}(\mu_I)$. In this section, we will show how to generalize the coupling in Section 5.1.1 to establish Lemma 5.

Lemma 5. Consider a K server system serving two-phase jobs with unequal rates. Consider any policy A and let the policies IF and A start from the same initial conditions and have the same arrival time process. Then there exists a coupling between IF and A such that $I_{\text{IF}}(t) \geq I_A(t)$ and $C_{\text{IF}}(t) \geq C_A(t)$ for all $t \geq 0$, where $I_{\text{IF}}(t)$ (resp. $I_A(t)$) is the number of inelastic transitions by time t under IF (resp. A), and $C_{\text{IF}}(t)$ (resp. $C_A(t)$) is the number of jobs completed by time t under IF (resp. A).

It is not trivial to apply the coupling in Section 5.1.1 to the situation where different phases (elastic and inelastic) have different exponential rates (μ_E and μ_I). The key component of our coupling in Section 5.1.1 was that we could parse time into blocks of length $\text{Exp}(K\mu)$ to keep both systems in sync. Now, the size of the blocks could depend on which types of phases (elastic or inelastic) are being served, and thus may be unequal between the two systems.

To tackle this problem, we leverage the technique of Markov chain *uniformization*. In uniformization, we find a rate μ that is at least as large as any transition rate in the Markov chain. Our goal is to set the total transition rate out of both the elastic and inelastic states to be μ without changing the behavior of the Markov chain. For instance, if $\mu_E > \mu_I$, we take $\mu := \mu_E$. Clearly, the total transition rate out of the elastic state is already μ in this case. For the inelastic state, because $\mu_I < \mu$, we add a self-loop transition that brings the total transition rate out of the state up to μ . We can then rewrite both μ_I and the new self-loop transition in terms of μ . Specifically, we can think of first transitioning out of the inelastic state with some total rate μ , and then either moving to the completion state or experiencing a self-loop with probability p_I or $1 - p_I$ respectively. By choosing $p_I = \frac{\mu_I}{\mu}$, we can re-write μ_I as $p_I \mu$. We then write the self-loop transition rate as $(1 - p_I)\mu$. Fig. 5(a) shows the resulting uniformized Markov chain.

It is easy to confirm that, since we have only added some self-loop transitions, the uniformized Markov chain is equivalent to the original Markov chain (Fig. 4(a)). The case where $\mu_E < \mu_I$ is symmetric and the uniformized Markov chain for this case is shown in Fig. 5(b).

To present our coupling, we will use the uniformized job model. Under the uniformized job model, $I_A(t)$ can differ from the total number of job completions, $C_A(t)$. However, for the coupling we present, we can use the number of inelastic transitions to directly recover the total number of job completions.

5.2.1. System coupling

Our goal in the coupling is twofold. Once again, we want to chop up time into blocks of length $\text{Exp}(K\mu)$ to keep S_{IF} and S_A roughly in sync. Additionally, we want to construct a coupling where reasoning about the number of inelastic transitions, $I_A(t)$, suffices for reasoning about total job completions, $C_A(t)$. That is, we want to find a coupling under which $I_{\text{IF}}(t) \geq I_A(t), \forall t \geq 0$ implies $C_{\text{IF}}(t) \geq C_A(t), \forall t \geq 0$.

Job arrivals: S_{IF} and S_A share the same arrival time sequence, and start with the same initial conditions.

Job transitions and departures: Our coupling in this case closely follows the coupling in Section 5.1.1. Specifically, the current time t is updated in the same manner as in Section 5.1.1. However, we handle potential transitions slightly differently, due to uniformization. We only discuss the case $\mu_I < \mu_E$, as the reverse case can be handled symmetrically.

When $\mu_I < \mu_E$, we take our dominating rate to be $\mu := \mu_E$. We generate an infinite sequence, $(X_n)_{n \geq 1}$, of i.i.d. $\text{Bern}(p_I)$ random variables, where $p_I = \frac{\mu_I}{\mu}$. The realizations of $(X_n)_{n \geq 1}$ are shared between the two systems. These *coin flips* will determine whether an inelastic transition results in a self-loop or in a job completion.

Throughout time, both systems keep track of the total number of inelastic transitions which have occurred. More concretely, each system starts with its own counter, n , which is initialized to 0. For each system, if we randomly select a server holding an inelastic job while experiencing a potential transition, we increment this system's counter ($n \leftarrow n + 1$).

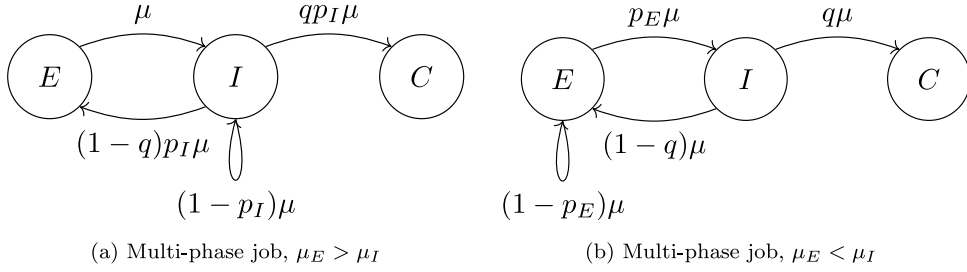


Fig. 6. Two cases of uniformizing multi-phase jobs. In Fig. 6(a), $\mu_E > \mu_I$, so we take our dominating rate as $\mu := \mu_E$. We then take $p_I := \frac{\mu_I}{\mu}$, and set the total transition rate out of the inelastic state to be μ . With probability $1 - p_I$, after completing an inelastic phase, we immediately start another one. With complementary probability p_I , the job does one of two things. With probability q , it completes. Otherwise, with probability $1 - q$, it begins an elastic phase. Fig. 6(b) shows the analogous case where $\mu_E < \mu_I$.

We then check position n of the shared infinite sequence of coin flips. If $X_n = 1$, the inelastic job completes and exits the system. Otherwise, we have a self-loop transition and no job exits the system. Hence, for any policy, π , at time t , we have

$$C_\pi(t) = \sum_{i=1}^{I_\pi(t)} X_i.$$

Since S_{IF} and S_A share a common sequence of coin flips, $I_{IF}(t) \geq I_A(t)$ implies $C_{IF}(t) \geq C_A(t)$.

5.2.2. Proof of Lemma 5

Since we only need to show $I_{IF}(t) \geq I_A(t)$, $\forall t \geq 0$, we can use the proof of Lemma 4 found in the Appendix verbatim to prove Lemma 5.

6. Optimality in the general case

We now consider the fully general multi-phase job structure, as seen in Fig. 4(c). In order to prove Theorem 1, it suffices to prove Lemma 6.

Lemma 6. Consider a K server system serving multi-phase jobs. Consider any policy A and let the policies IF and A start from the same initial conditions and have the same arrival time process. Then there exists a coupling between IF and A such that $I_{IF}(t) \geq I_A(t)$ and $C_{IF}(t) \geq C_A(t)$ for all $t \geq 0$, where $I_{IF}(t)$ (resp. $I_A(t)$) is the number of inelastic transitions by time t under IF (resp. A), and $C_{IF}(t)$ (resp. $C_A(t)$) is the number of jobs completed by time t under IF (resp. A).

As in Section 5.2, we use uniformization to rewrite the job structure of Fig. 4(c) so that the elastic and inelastic phase transitions have equal rates. There are two possible uniformizations here, once again depending on how μ_E and μ_I relate. Determining the dominating rate μ and transition probabilities p_I or p_E is the same as in Section 5.2, and the two possible uniformized Markov chains are shown in Fig. 6. With these job structures in mind, we present the system coupling which allows us to prove the optimality of IF.

6.1. System coupling

As in Section 5.1.1, we construct a coupling that keeps systems S_A and S_{IF} in sync with respect to potential transition times and that allows us to use $I_A(t)$ to reason about $C_A(t)$.

Job arrivals: As in Sections 5.1.1 and 5.2.1, we let S_{IF} and S_A share the same arrival time sequence and start with the same initial conditions.

Job transitions and departures: For the most part, the transition process is similar to the uniformized case presented in Section 5.2.1. However, while we previously only needed a single infinite sequence of i.i.d. Bernoulli random variables, here we will need two. We state the two cases ($\mu_E < \mu_I$ and $\mu_E > \mu_I$) separately, as they differ slightly in their construction.

First, we consider $\mu_E < \mu_I$ (Fig. 6(b)). Here, instead of a single sequence of coin flips, we have two shared sequences of coin flips. The first sequence, $(X_n)_{n \geq 1}$, is an i.i.d. sequence of $Bern(p_E)$ random variables. If $X_n = 1$, the n th elastic transition results in an elastic phase completion, producing an inelastic phase. Otherwise, if $X_n = 0$, the elastic transition does not result in a phase completion. The second sequence, $(Y_n)_{n \geq 1}$, is a sequence of i.i.d. $Bern(q)$ random variables. Recall that q is the probability that the completion of an inelastic phase will result in a job completion. If $Y_n = 0$, the n th inelastic transition results in the creation of an elastic phase. If $Y_n = 1$, the n th inelastic transition results in a job completion.

The case when $\mu_E > \mu_I$ (Fig. 6(a)) is slightly more complex. Here, we do not have any self-loops for elastic phases. However, there are three possible outcomes for inelastic phases. We therefore keep track of two sequences of i.i.d.

Bernoulli random variables, $(X_n)_{n \geq 1}$ and $(Y_n)_{n \geq 1}$. In the first sequence, X_n is distributed as $Bern(p_I)$. In the second sequence, Y_n is distributed as $Bern(q)$.

If $X_n = 0$, the n th inelastic transition does not result in the completion of an inelastic phase. If $X_n = 1$, the n th inelastic transition results in a phase completion, and we then examine the sequence (Y_n) . If the n th inelastic transition results in the m th overall inelastic phase completion, we check Y_m . If $Y_m = 1$, the job completes, and if $Y_m = 0$, the job transitions to an elastic phase.

Because S_{IF} and S_A share the same sequence of coin flips, comparing the number of inelastic transitions between systems is equivalent to comparing the number of job completions. That is, if $I_{IF}(t) \geq I_A(t)$, $\forall t \geq 0$, then $C_{IF}(t) \geq C_A(t)$.

6.2. Proof of Lemma 6

Multi-phase jobs add an extra layer of complexity which prevents us from directly leveraging the arguments used in Lemmas 4 and 5. When an inelastic phase completes, there are two possible outcomes: either an elastic phase will be produced or a job will complete. Our insight is that we can view the creation of an elastic phase as a job completion immediately followed by an arrival of a job in an elastic phase. This reduction puts us back in the case of two-phase jobs with unequal rates, allowing us to invoke Lemma 5. We formalize this argument in the proof of Lemma 6.

Proof of Lemma 6. Consider the above coupling under any given sample path. First, we replace each inelastic transition that produces an elastic phase with a different type of transition. Namely, we replace these transitions with a job completion followed immediately by an arrival of a job in an elastic phase. We will refer to this replacement as our *re-framing* of the sample path. Observe that the schedules produced in S_{IF} and S_A remain the same under the re-framing. While the number of job completions by any point t , $C_{IF}(t)$ and $C_A(t)$, may change under this re-framing, the key insight is that the number of inelastic transitions, $I_{IF}(t)$ and $I_A(t)$ respectively, remains identical. Thus, if we can argue that IF maximizes the number of inelastic transitions by any point in time under the re-framing, it does so in the original environment as well. This is sufficient for proving that $C_{IF}(t) \geq C_A(t)$, $\forall t \geq 0$ under the original sample path.

Second, observe that our proof of Lemma 5 still holds if we allow additional arrivals at potential transition times, so long as these arrivals occur simultaneously in both systems. However, under our re-framing, the arrivals we add may not occur simultaneously in S_{IF} and S_A since they are generated by inelastic transitions to elastic phases. We address this issue by establishing the following claim.

Claim. Let the sequence of additional arrival times under the re-framing be (t_n) in S_{IF} and (s_n) in S_A . For any $n \geq 1$, we have $t_n \leq s_n$, i.e. the n th additional arrival occurs in S_{IF} before it occurs in S_A .

We will prove this claim below, allowing us to complete the proof of Lemma 6. Specifically, for any time t , let n be the index such that $t_n \leq t < t_{n+1}$. The claim tells us that S_A experiences additional arrivals at $s_1 \geq t_1, s_2 \geq t_2, \dots, s_{n+1} \geq t_{n+1}$. However, we can view S_A as a system that has additional arrivals at $t_1, t_2, \dots, t_{n+1}$, but chooses to not schedule these additional arrivals until after $s_1, s_2, \dots, s_{n+1}$. Then by Lemma 5, we have $I_{IF}(t) \geq I_A(t)$, which completes the proof of Lemma 6. $\square$

Proof of the claim. We will show inductively that the n th additional arrival occurs in S_{IF} before it does in S_A . We first argue that $t_1 \leq s_1$. Observe that, on the time interval $[0, t_1 \wedge s_1]$, S_{IF} and S_A experience precisely the same sequence of arrivals. Hence, by Lemma 5, IF maximizes the number of inelastic transitions by any time $t \in [0, t_1 \wedge s_1]$. In particular, it maximizes the number inelastic transitions by time $t_1 \wedge s_1$. Since S_{IF} and S_A share the same sequences of Bernoulli random variables, it must be that the system with more inelastic transitions experiences the first inelastic to elastic transition, and hence $t_1 \leq s_1$. Now note that the schedule produced by S_A is identical to that produced by a policy which receives the additional arrival at time t_1 instead of time s_1 , but just chooses to ignore its existence until later on. This allows us to assume that the extra arrival into S_A occurs at t_1 instead of s_1 . We then observe that, on the interval $[0, s_2 \wedge t_2]$, if systems S_{IF} and S_A experience the same sequence of arrivals then IF maximizes the number of inelastic transitions by time $s_2 \wedge t_2$. Hence, we have that $t_2 \leq s_2$.

Using the same iterative argument, it follows that S_A and S_{IF} can be assumed to experience the same sequence of arrivals up to time $s_n \wedge t_n$, and thus by Lemma 5 we have that $t_n \leq s_n$. $\square$

Having proven Lemma 6, we can now prove Theorem 1.

Theorem 1. Consider a K server system serving multi-phase jobs. The policy IF stochastically maximizes the number of jobs completed by any point in time. Specifically, for a policy A , let $C_A(t)$ denote the number of jobs completed by time t and let $N_A(t)$ denote the number of jobs in the system at t . Then under any arbitrary arrival time process, $C_{IF}(t) \geq_{st} C_A(t)$ for all times $t \geq 0$. Consequently, $N_{IF}(t) \leq_{st} N_A(t)$ for all times $t \geq 0$.

Proof. Lemma 6 implies the existence of a coupling such that $C_{IF}(t) \geq C_A(t)$, $\forall t \geq 0$. Consequently, $C_{IF}(t) \geq_{st} C_A(t)$, $\forall t \geq 0$. Since the number of jobs in the system at time t is just the total number of arrivals by time t minus the total number of completions by time t , the claim $N_{IF}(t) \leq_{st} N_A(t)$, $\forall t \geq 0$ also readily follows from Lemma 6. $\square$

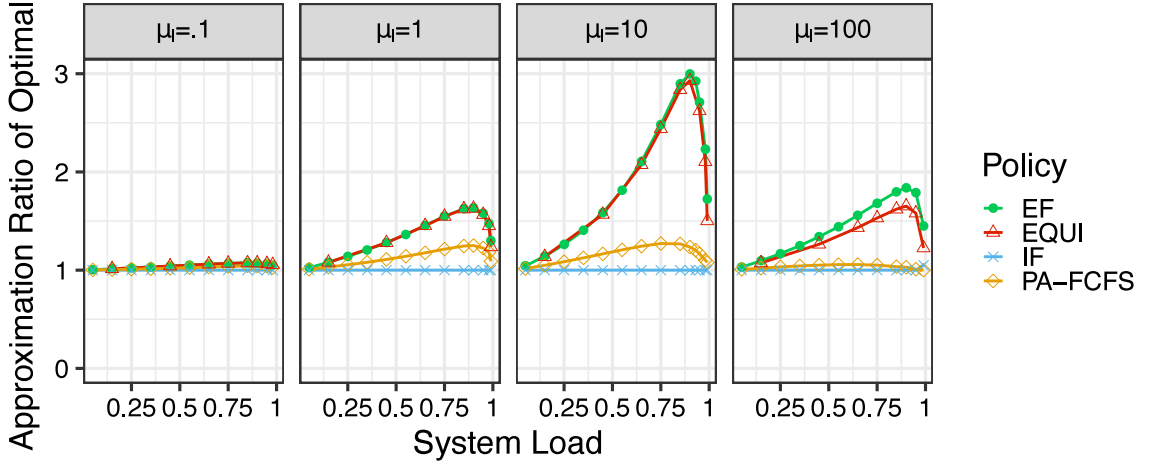


Fig. 7. The approximation ratio of mean response time under EQUI, EF, PA-FCFS, and IF when compared with the optimal mean response time. Phases have exponentially distributed inherent sizes. IF is optimal (see Section 6) and thus has an approximation ratio of 1. In each case, $K = 100$, $\mu_E = 1$, $q = 0.2$, and jobs arrive according to a Poisson process. All jobs begin with an elastic phase. Results are shown as μ_I varies from $\mu_I = 0.1$ (the rare case where inelastic phases are long compared to elastic phases) to $\mu_I = 100$ (the more common case where inelastic phases are short compared to elastic phases).

7. Evaluation

The analysis of Section 6 has shown that, when jobs have the structure presented in Fig. 2, IF is optimal. In particular, IF minimizes the steady-state mean response time for any settings of μ_E , μ_I , q , and any arrival time process such that the system is stable.

The purpose of this section is three-fold. First, we examine the benefit of doing IF as opposed to other scheduling policies used in real-world systems or proposed in the literature. Second, we will relax the assumption that the phases are exponentially distributed and consider a range of phase size distributions from low-variability to high-variability. Third, we will consider jobs that consist of a deterministic number of phases and may end with an elastic phase, relaxing our assumption that jobs follow the underlying structure illustrated in Fig. 2. We find that, even with these relaxed assumptions, IF is almost always a great choice compared to all other competitor policies.

We begin by describing the competitor policies in Section 7.1. Then we show the comparisons to IF via simulation in Sections 7.2 and 7.3.

7.1. Competitor scheduling policies

We compare IF to three competitor policies.

EQUI is a phase-unaware policy for scheduling parallelizable jobs that has been widely advocated for in both the worst-case [15,19,20] and stochastic [18] theoretical literature. EQUI divides servers equally among all jobs in the system. That is, if there are N jobs in the system, each job receives an allocation of $\frac{K}{N}$ servers regardless of its current phase.

Phase-Aware First-Come-First-Served (PA-FCFS) is a popular policy in systems applications because it is easy to implement with little space or time overhead. PA-FCFS proceeds by iteratively looking at the next earliest arriving job in the system and allocating as many servers as possible to this job until all servers have been allocated. (A job in an inelastic phase is obviously allocated only 1 server.)

Elastic-First (EF) gives strict priority to the earliest arriving job in an elastic phase. If no jobs are in an elastic phase, servers are allocated to any jobs in an inelastic phase in FCFS order. Intuitively, EF seems like it might perform well in cases where elastic phases are smaller than inelastic phases on average. In this case, EF can be thought of as a greedy policy which continuously minimizes the expected time until the next phase completion. This is analogous to the GREEDY* policy proposed in [18]. However, we will see that this intuition is wrong.

7.2. Evaluation of policies under our job model

Fig. 7 shows the results of simulations comparing the performance of IF, EQUI, PA-FCFS, and EF under the model defined in Section 3 as we vary μ_I . Each simulation consists of 100 million job completions. Although we have already proven the optimality of IF in these cases, Fig. 7 shows that the improvement of IF over the competitor policies is significant. In this small sample of the parameter space, IF outperforms PA-FCFS by up to 25%, and outperforms EF and EQUI by as much as a factor of 3. It is interesting to note that IF outperforms EF even when $\mu_E > \mu_I$. Even in this case, EF suffers from its failure to defer parallelizable work.

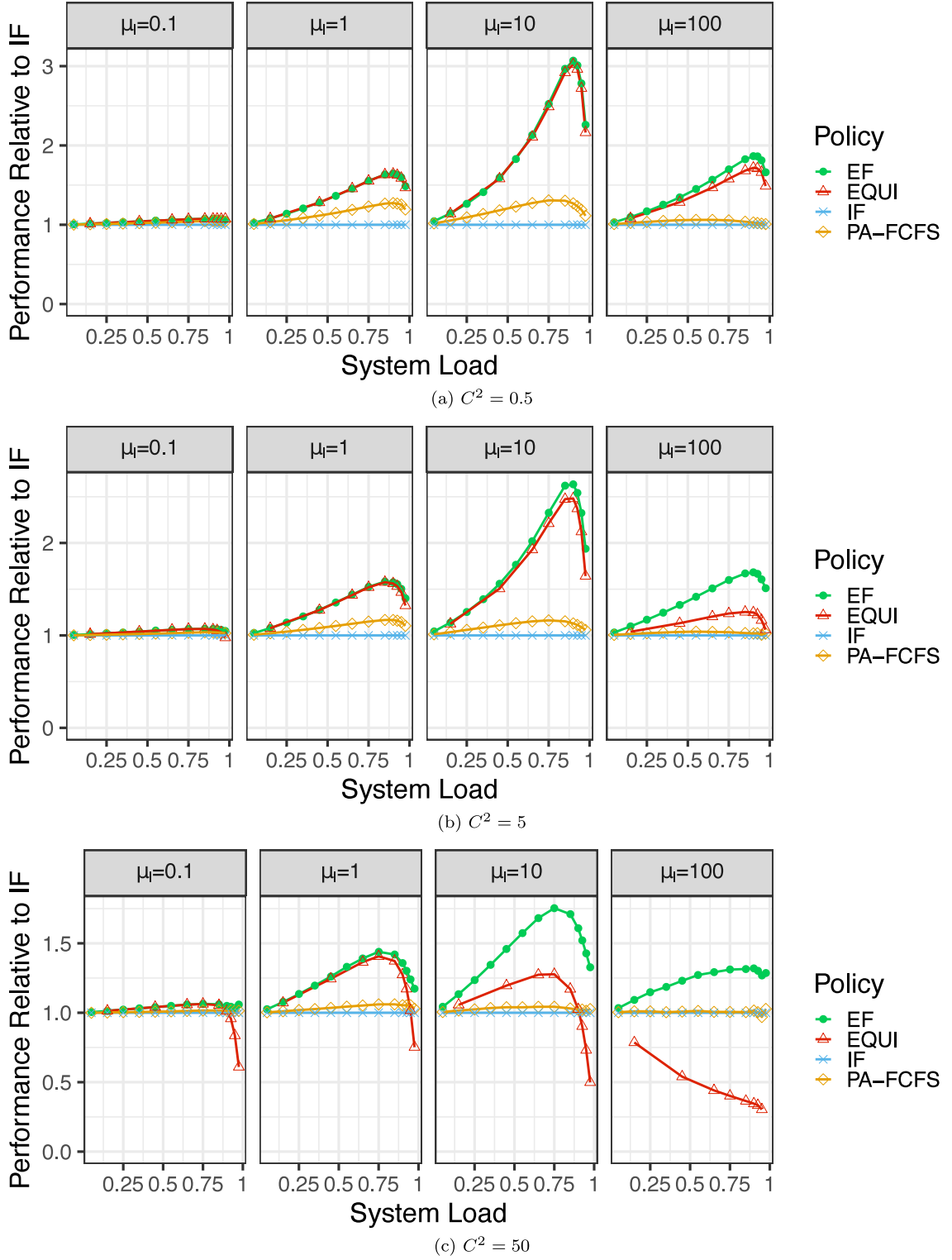


Fig. 8. The approximation ratio of mean response time under EQUI, EF, PA-FCFS, and IF, all compared with IF, when phases follow a Weibull distribution. In each case, $K = 100$, $\mu_E = 1$, $q = 0.2$, and jobs arrive according to a Poisson process. IF typically still outperforms the competitor policies. When jobs are highly variable ($C^2 = 50$), EQUI outperforms IF due to its insensitivity to job size variance.

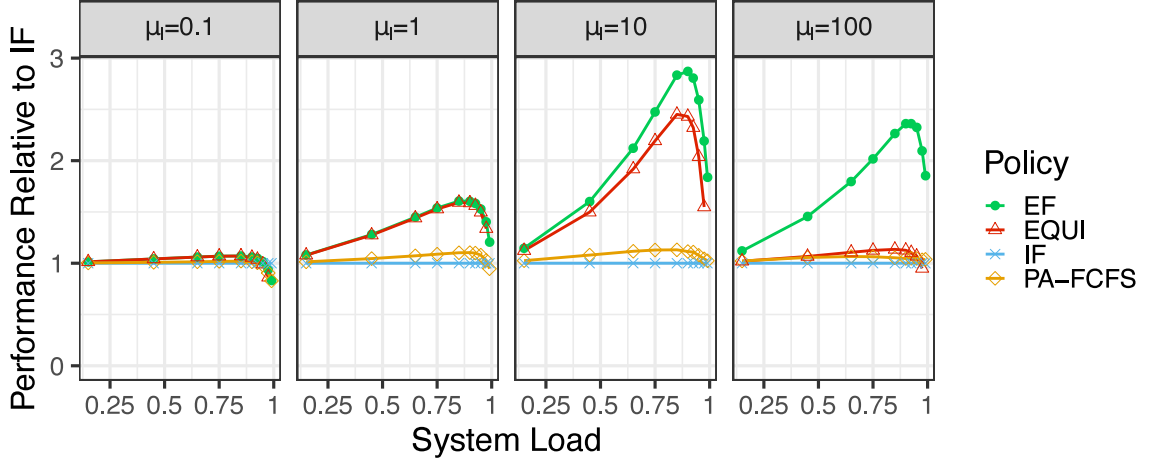


Fig. 9. The approximation ratio of mean response time under EQUI, EF, PA-FCFS, and IF, all compared with IF, when jobs consist of a deterministic number of phases. In each case, $K = 100$, $\mu_E = 1$, and jobs arrive according to a Poisson process. Each job consists of 5 phases whose sizes are Weibull distributed with $C^2 = 5$. Each phase, including the last phase, is chosen to be either elastic or inelastic with probability .5. Although these changes to the job structure violate our theoretical assumptions, IF is still generally the best of the policies we consider.

7.3. Jobs with alternate phase size distributions

Although we have shown that IF is optimal when phase sizes are exponentially distributed, we wish to further show that IF outperforms other policies under a range of phase size distributions. To examine the sensitivity of IF's performance to the underlying phase size distributions, we examine different distributions with a range of variances. Specifically, we consider the case where phases are Weibull distributed and the *squared coefficient of variation*, C^2 , of the phase size distribution is both higher and lower than that of an exponential distribution.⁴

Fig. 8 shows the performance of each competitor policy in simulation relative to the performance of IF (hence, the performance of IF is always normalized to 1). In most cases, IF is still the best of the four policies by a wide margin.

When $C^2 = 50$ we do find examples where EQUI outperforms IF. Here, when job sizes are highly variable, EQUI benefits from its insensitivity to the variance of the job size distribution [18]. Specifically, because the phase size distributions have decreasing failure rates, working on phases with the *least attained service* will generally result in completing smaller phases before larger phases [37]. EQUI biases in this direction by dividing servers equally amongst all jobs in the system.

The relative performance of the competitor policies compared to IF also depends on the distribution of the number of phases comprising each job. For instance, when $q = 1$ and all jobs begin with an elastic phase, IF and PA-FCFS are equivalent policies. However, as q decreases, for a given system load, the gap between IF and PA-FCFS widens as it becomes increasingly likely that PA-FCFS will make a mistake and give priority to an elastic phase. Hence, although Fig. 8 shows that EQUI can outperform IF when $q = .2$, in the case where $q = .025$ IF again outperforms EQUI at all loads.

7.4. Jobs with alternate structures

We can also examine the effect of the job structure on the performance of IF by considering cases where the sequence of a job's phases is not determined by a Markov chain of the form shown in Fig. 2. Specifically we consider the case where the total number of job phases is deterministic. Furthermore, we allow each of the job phases to be either elastic or inelastic, including the final phase of each job. The results of these experiments look largely the same as the results in Fig. 8, and an example of these results is shown in Fig. 9. Although IF is not optimal in this case, it is still frequently the best of the policies we consider.

One notable exception where IF becomes worse than the competitor policies occurs when $\mu_l = 0.1$ and system load is high. In this case, the benefit of doing IF has been diminished in several ways. First, the benefit of deferring parallelizable work is realized when there are very few jobs in the system, but the scheduling policy is still able to utilize all K servers because at least one job is in an elastic phase. Under very high load, the probability of having few jobs in the system vanishes, decreasing the benefit of deferring parallelizable work. Second, the change in job structure has made it possible for IF to work on jobs with longer expected remaining total sizes. This is an effect of both allowing jobs to end in elastic phases and making the total number of phases for each job deterministic. Finally, we are considering the case where

⁴ A Weibull distribution with shape parameter $k = 1$ collapses to an exponential distribution ($C^2 = 1$). Adjusting k changes the distribution to have either higher C^2 ($k < 1$) or lower C^2 ($k > 1$) than an exponential distribution.

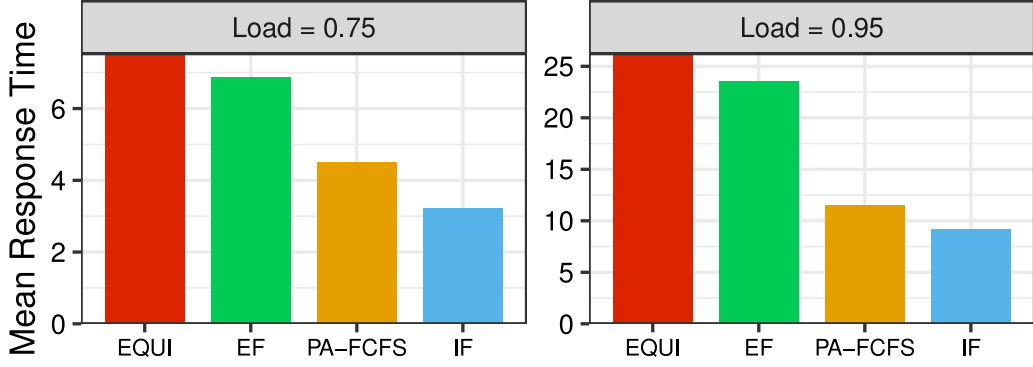


Fig. 10. The mean response time of EQUI, EF, PA-FCFS, and IF processing a workload consisting of a mixture of 5 queries from the Star Schema Benchmark. We assume Poisson arrivals. Although this workload violates our modeling assumptions, IF is still the best policy by a wide margin. IF improves upon the next best policy, the PA-FCFS policy used in the NoisePage database, by up to 30%.

$\mu_I = 0.1$, meaning inelastic phases are ten times larger than elastic phases on average. Because IF prioritizes inelastic phases, this particular choice of the phase size distributions further increases the chance that IF suffers by completing jobs with larger total sizes before jobs with smaller total sizes. Despite all of these challenges, IF still performs within 18% of the best competitor policy.

8. Case study: Scheduling in databases

Throughout this paper, we have drawn inspiration for our model from a range of systems including modern databases. In this section, we consider whether the scheduling policies that we have proposed work well for real database workloads. Because this section specifically considers scheduling in databases, we will refer to a scheduling policy as allocating *cores* to *queries* rather than allocating servers to jobs. Real database workloads differ from our modeling assumptions in two ways. First, phase sizes are not exponentially distributed. Second, the sequence of phases for each query is not determined by an underlying Markov chain. In this case study, we ask whether IF will still perform well under these real-world conditions.

To answer this question, we perform simulations using a workload consisting of a mixture of five queries from the Star Schema Benchmark. The ordering of phases and the phase sizes in our simulations are based on the timings of actual queries running in the NoisePage database, as shown in Fig. 1. Each query consists of a deterministic number of elastic and inelastic phases (either six or seven total phases per query). Of the five queries we consider, four queries end with an inelastic phase and one query ends with an elastic phase. For each arrival in our simulation, a query type is drawn uniformly at random. The query sizes have a squared coefficient of variation of $C^2 = 0.41$, meaning that query sizes are not highly variable in this case. In particular, note that these queries clearly deviate from the job structure illustrated in Fig. 2.

Fig. 10 shows the results of these simulations. The ordering of the policies with respect to mean response time is the same as what we observed in Section 7. In particular, IF is again consistently the best of the policies we consider, and IF outperforms the PA-FCFS policy used in the current version of NoisePage by up to 30%.

8.1. Size-aware scheduling

Although the focus of this paper has been the setting where job sizes are unknown to the scheduler, we recognize that schedulers in real-world databases often have knowledge of the size of each query phase and the number of phases comprising each query. Specifically, the query planner in the NoisePage database on which we have based our simulations can provide the scheduler with information about the sequence of phases for each query and an estimate of phase sizes in addition to information about the current phase.

Historically, when job sizes are known, the performance modeling community has advocated for reducing mean response time by trying to complete smaller jobs before larger jobs [38,39]. This begs the question of whether the phase-aware policies developed in this paper can be improved by adapting them to favor short jobs. Notably, NoisePage and many other databases use a PA-FCFS policy, and do not leverage the available information about query sizes to make better scheduling decisions. Would favoring short queries improve response times in these systems?

Our theorems in Sections 5 and 6 have shown the importance of *deferring parallelizable work* by giving priority to inelastic phases in order to maintain the overall efficiency of the system. In the case where phase sizes are known, it is not immediately clear how to balance the objectives of favoring shorter queries and deferring parallelizable work.

Size-aware Scheduling Policies. We now consider two size-aware scheduling policies that favor queries with smaller *remaining total size*, the sum of the remaining sizes of all of a query's remaining phases. As we will see, one of the

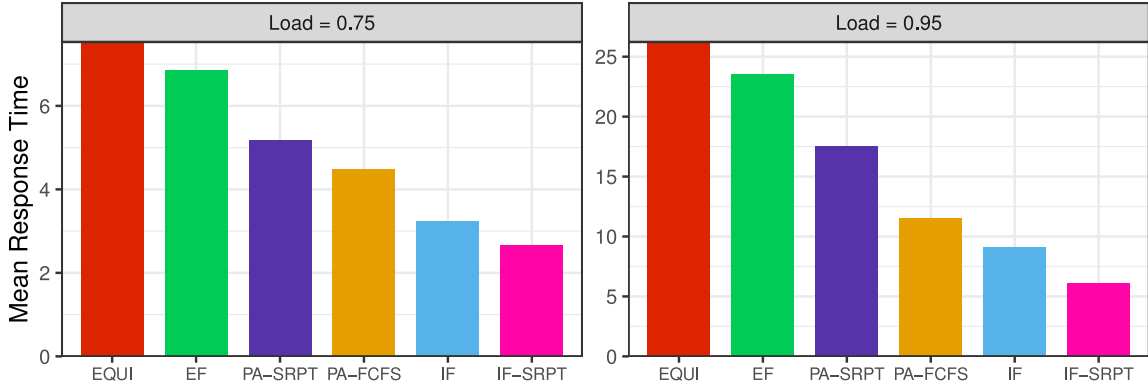


Fig. 11. Comparison with size-aware policies. The mean response time of EQUI, EF, PA-SRPT, PA-FCFS, IF and IF-SRPT processing a workload consisting of a mixture of 5 queries from the Star Schema Benchmark. We assume Poisson arrivals. IF-SRPT can improve upon IF by 33% by leveraging query size information. Notably, PA-SRPT performs worse than PA-FCFS despite attempting to leverage size information.

scheduling policies performs well because it manages to both favor short queries and grant strict priority to inelastic phases. However, the other policy, which favors the shortest queries in the system but does not otherwise defer parallelizable work, does even worse than PA-FCFS.

The first policy we consider is an adaptation of IF to the case where query sizes are known to the scheduler. We call this new policy *Inelastic-First-Shortest-Remaining-Processing-Time* (IF-SRPT) because it combines IF with the ubiquitous SRPT scheduling policy. IF-SRPT gives strict priority to inelastic phases over elastic phases in the same manner as IF. However, among inelastic phases, IF-SRPT gives priority to the phases belonging to the queries with the smallest remaining total sizes. Likewise, when choosing to run an elastic phase, IF-SRPT will choose the elastic phase belonging to the query with the smallest remaining total size.

Our second policy is a *Phase-Aware SRPT* policy, which we refer to as PA-SRPT. PA-SRPT gives strict priority to the phases belonging to the queries with the smallest remaining total sizes, regardless of whether a phase is elastic or inelastic. However, PA-SRPT is phase-aware in that it avoids allocating too many cores to inelastic phases. Hence, if the query with the smallest remaining total size is in an inelastic phase, PA-SRPT will allocate one core to this query. If the next smallest query is in an elastic phase, PA-SRPT will allocate the remaining $K - 1$ cores to this second smallest query. Although PA-SRPT does not explicitly defer parallelizable work, it biases more strongly towards the shortest queries in the system than IF-SRPT does.

We again evaluate these policies using a workload based on the Star Schema Benchmark, and the results are shown in Fig. 11. Unsurprisingly, IF-SRPT is the best performer. It benefits from biasing its allocations towards shorter queries while still deferring parallelizable work. This leads IF-SRPT to achieve a mean response time which can be 33% lower than that of IF, and 47% lower than that of the PA-FCFS policy used in NoisePage. What is more counterintuitive is that PA-SRPT performs quite poorly. In fact, PA-SRPT is worse than PA-FCFS in both cases shown in Fig. 11, and IF-SRPT outperforms PA-SRPT by up to a factor of 3.

8.2. Why PA-SRPT is worse than PA-FCFS

As seen in this paper, deferring parallelizable work is vital to reducing mean response time. PA-SRPT suffers from its failure to defer parallelizable work. It is not immediately clear, however, why PA-SRPT is even worse than PA-FCFS, given that neither policy explicitly defers parallelizable work.

Although neither PA-SRPT nor PA-FCFS explicitly defers parallelizable work, we can see that PA-SRPT suffers because it inadvertently defers far less parallelizable work than PA-FCFS. We define the *percentage of deferred parallelizable work* under a given policy at time t to be the number of cores allocated to inelastic phases divided by the number of cores that IF would allocate to inelastic phases. We can then consider the long-run time-average percentage of deferred parallelizable work under various policies. We normalize this quantity using the allocations under IF because IF allocates as many cores to inelastic phases as possible without being wasteful. As a result, IF defers 100% of parallelizable work by definition. Phase-unaware policies, such as EQUI, can defer more than 100% of parallelizable work by wastefully allocating too many cores to inelastic phases.

Fig. 12 shows that PA-SRPT defers far less parallelizable work than PA-FCFS, leading PA-SRPT to perform poorly. The poor performance of PA-SRPT is due to the structure of one particular query, which is both small in total size and ends with an elastic phase. PA-SRPT tends to give strict priority to this type of query, while PA-FCFS treats all queries equally. This leads PA-SRPT to defer less parallelizable work than PA-FCFS. Fig. 12 also shows that, while IF-SRPT manages to favor short jobs, it still defers parallelizable work. IF-SRPT is able to both defer 100% of parallelizable work and prioritize shorter queries, leading to lower mean response time.

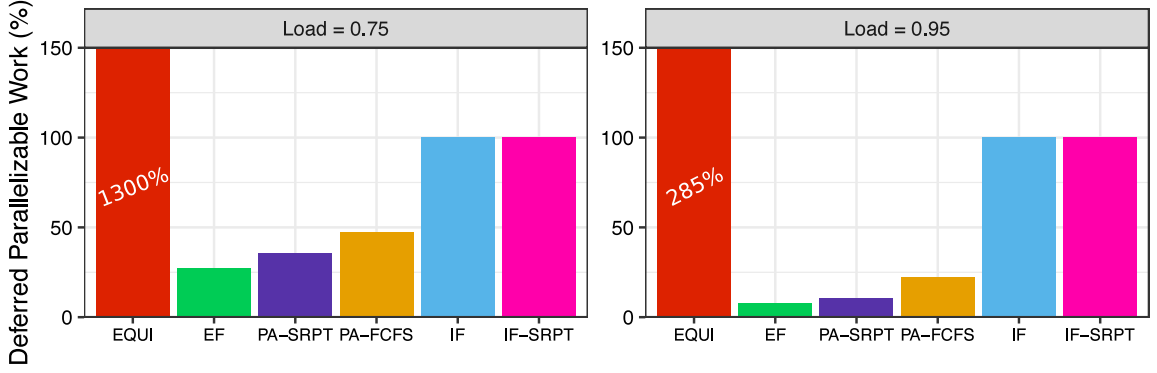


Fig. 12. The percentage of deferred parallelizable work under EQUI, EF, PA-SRPT, PA-FCFS, IF and IF-SRPT given a workload consisting of a mixture of 5 queries from the Star Schema Benchmark. IF-SRPT defers 100% of parallelizable work, but PA-SRPT defers even less parallelizable work than PA-FCFS.

9. Conclusion

This paper addresses the optimal scheduling of parallelizable jobs, specifically jobs that consist of different numbers of elastic and inelastic phases. While optimality results in the literature often involve asymptotic approximations such as scaling of system size or heavy traffic assumptions, the results in this paper make no such assumptions. We prove that the IF policy, which defers parallelizable work, is optimal in a strong sense; for any number of servers, K , for any system load, ρ , for any arrival process (including adversarial arrivals), and when jobs can each consist of an arbitrary number of phases. While our proofs do require that the phases have exponentially distributed sizes and that jobs all end with inelastic phases, experimental evaluation shows that the dominance of IF typically extends to cases when job sizes and structures deviate from these assumptions. Furthermore, IF does not need knowledge of the job structure (other than knowing the current phase), i.e., IF does not require knowledge of the job parameters, μ_I , μ_E , and q .

We also show that IF performs well in simulation under database workloads where job structures, phase sizes, and phase types can vary widely. In the setting of scheduling in databases, it is common that the scheduler not only knows the phase of each job but also has knowledge of the job's size. Given that job sizes are known, it is natural to consider scheduling policies which favor short jobs. We consider two such policies: the PA-SRPT policy which strictly favors jobs with the shortest remaining total sizes and the IF-SRPT policy which both defers parallelizable work and favors short jobs. We find that IF-SRPT is far superior to PA-SRPT and performs even better than IF. This somewhat counterintuitive result underscores the importance of deferring parallelizable work when scheduling parallelizable jobs composed of phases.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgments

This work was supported by a Facebook Graduate Fellowship, a 2020 Google Faculty Award, a Google Research Award, an Infor Research Award, a Carnegie Bosch Junior Faculty Chair, and NSF grants NSF-CMMI-1938909, NSF-CSR-1763701, NSF-CCF-1824303, NSF-CCF-1845146, NSF-CCF-1733873, NSF-CNS-2007733.

Appendix. Proof of Lemma 4

Lemma 4. Consider a K server system serving two-phase jobs with equal rates. Consider any policy A and let the policies IF and A start from the same initial conditions and have the same arrival time process. Then there exists a coupling between IF and A such that $I_{IF}(t) \geq I_A(t)$ and $C_{IF}(t) \geq C_A(t)$ for all $t \geq 0$, where $I_{IF}(t)$ (resp. $I_A(t)$) is the number of inelastic transitions by time t under IF (resp. A), and $C_{IF}(t)$ (resp. $C_A(t)$) is the number of jobs completed by time t under IF (resp. A).

Proof. We proceed by induction on event times, as defined in Section 5.1.1. We parse time into two types of periods: *busy periods* where S_{IF} utilizes all K of its servers, and *idle periods* where S_{IF} idles at least one of its servers at any point in time. Let t_0 be the start of a period (either busy or idle). We show below that, if $I_{IF}(t_0) \geq I_A(t_0)$, then, at any point of

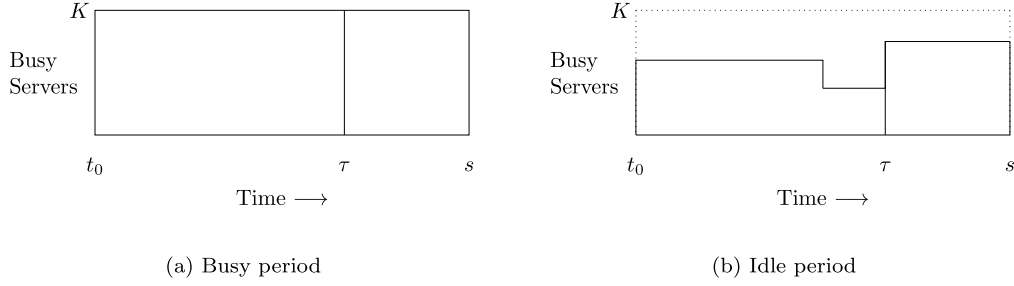


Fig. 13. Examples of busy and idle periods used in the proof of Lemma 4. The figures both show the relative positions of the times t_0 , τ , and s . The value K refers to the number of servers, and the heights of the boxes indicate how many servers S_{IF} allocates to jobs.

time t during the current period, $I_{IF}(t) \geq I_A(t)$. This claim is sufficient for proving Lemma 4 since time can be partitioned into disjoint alternating busy and idle periods. To get a sense of how time is partitioned, refer to Fig. 13.

As a base case for our induction, observe that $I_{IF}(0) = I_A(0)$. Depending on the initial conditions, time $t = 0$ will serve as either the start of the first busy period or the first idle period.

Busy Periods: We first consider the case where time t_0 marks the start of a busy period, and assume inductively that $I_{IF}(t_0) \geq I_A(t_0)$. We show that $I_{IF}(t) \geq I_A(t)$ for all times t in the busy period by contradiction. Assume for contradiction that there is some earliest time s in the busy period such that $I_{IF}(s) < I_A(s)$.

First, we argue that

$$N_{IF}^E(t_0) \leq N_A^E(t_0). \quad (1)$$

If $t_0 = 0$, this follows directly from the shared initial conditions of S_{IF} and S_A . Now suppose $t_0 > 0$. Observe that, since t_0 marks the beginning of a busy period, immediately before time t_0 , all of the jobs in S_{IF} must be in the inelastic phase. That is, $N_{IF}^E(t_0-) = 0$, and thus $N_{IF}^E(t_0-) \leq N_A^E(t_0-)$. Lastly, the event that happens at t_0 can only be job arriving since t_0 is the start of a busy period. Since the arrival occurs simultaneously in both systems, it follows that $N_{IF}^E(t_0) \leq N_A^E(t_0)$, as desired.

Next, let τ be the time for the event preceding the event at s . We claim that

$$I_{IF}(\tau) = I_A(\tau), \text{ and } N_{IF}(\tau) = N_A(\tau), \quad (2)$$

where $N_{IF}(\tau) = N_A(\tau)$ follows from $I_{IF}(\tau) = I_A(\tau)$. This is because the number of inelastic transitions determines the number of job completions and both systems experience the same arrival sequence. The claim $I_{IF}(\tau) = I_A(\tau)$ is true since $I_{IF}(t)$ is non-decreasing, $I_A(\tau)$ can increase by at most 1 at time s , and s is the earliest time during the busy period for which $I_{IF}(s) < I_A(s)$. More specifically, we can conclude that at time s , S_{IF} experiences an elastic transition, whereas S_A experiences an inelastic transition. This holds because $I_{IF}(s) < I_A(s)$ and $I_{IF}(\tau) = I_A(\tau)$ if and only if $I_{IF}(s) = I_{IF}(\tau)$ and $I_A(s) = I_A(\tau) + 1$. This implies that S_A experiences an inelastic transition at time s . Furthermore, S_{IF} experiences an elastic transition at time s since IF does not idle servers during a busy period.

Now, we can claim that

$$E_{IF}(t_0, s) \leq E_A(t_0, s). \quad (3)$$

Since we have shown that $N_{IF}^E(t_0) \leq N_A^E(t_0)$ in (1), it suffices to show that $N_{IF}^E(s) \geq N_A^E(s)$. Per our coupling, the previous paragraph implies that S_{IF} is running fewer inelastic jobs on the interval $[\tau, s)$ than S_A . Since S_{IF} always runs the maximal number of inelastic jobs, we have that $N_{IF}^I(\tau) < N_A^I(\tau)$. Moreover, since $N_{IF}(\tau) = N_A(\tau)$ by (2), we know that $N_{IF}^E(\tau) > N_A^E(\tau)$, and thus $N_{IF}^E(s) \geq N_A^E(s)$.

Finally, let M denote the number of potential transitions during $(t_0, s]$. Since S_{IF} is never idling servers between times t_0 and s , we have the identities:

$$M = E_{IF}(t_0, s) + I_{IF}(t_0, s), \text{ and } M \geq E_A(t_0, s) + I_A(t_0, s).$$

Consequently, utilizing (3) and rearranging, we have that:

$$I_{IF}(t_0, s) \geq I_A(t_0, s). \quad (4)$$

Moreover, recall that by definition, $I_{IF}(s) = I_{IF}(t_0) + I_{IF}(t_0, s)$ and $I_A(s) = I_A(t_0) + I_A(t_0, s)$. Since we assumed $I_{IF}(t_0) \geq I_{IF}(t_0)$, and we know that $I_{IF}(t_0, s) \geq I_A(t_0, s)$ by (4), we have $I_{IF}(s) \geq I_A(s)$, a contradiction. This completes the induction step for busy periods.

Idle periods: Next, we consider the case that time t_0 marks the beginning of an idle period, and again inductively assume that $I_{IF}(t_0) \geq I_A(t_0)$. To show $I_{IF}(t) \geq I_A(t)$ for all times t in the idle period, we once again proceed by contradiction. That is, suppose there is some earliest time s in the period such that $I_{IF}(s) < I_A(s)$.

First, observe that, since S_{IF} always chooses to idle at least one server during the idle period, there cannot be any elastic phase jobs in the system. That is, $N_{IF}^E(t) = 0$ for all times t in the idle period.

Letting τ be defined again as the time for the event preceding the event at s , by a similar reasoning to before, we must have that

$$I_{IF}(\tau) = I_A(\tau), \quad (5)$$

and that, at time s , S_{IF} does not have a transition, whereas S_A experiences an inelastic transition.

Now we show that we have a contradiction. First note that the equality $I_{IF}(\tau) = I_A(\tau)$ in (5) implies that $N_{IF}(\tau) = N_A(\tau)$. Next, since at time s , S_{IF} does not have a transition but S_A experiences an inelastic transition, per our coupling, S_{IF} is running strictly fewer jobs in the inelastic phase than S_A . Since S_{IF} has no elastic jobs, this implies that $N_{IF}(\tau) < N_A(\tau)$, leading to a contradiction. This completes the induction step for idle periods. $\square$

References

- [1] NoisePage - The self-driving database management system, <https://noise.page>.
- [2] Patrick O'Neil, Elizabeth O'Neil, Xuedong Chen, Stephen Revilak, The star schema benchmark and augmented fact table indexing, in: Performance Evaluation and Benchmarking: First TPC Technology Conference, TPCTC 2009, Lyon, France, August 24–28, 2009, Revised Selected Papers, 2009, pp. 237–252.
- [3] Nathan R. Tallent, John M. Mellor-Crummey, Effective performance measurement and analysis of multithreaded applications, in: Proceedings of the 14th ACM SIGPLAN Symposium on Principles and Practice of Parallel Programming, 2009, pp. 229–240.
- [4] Thu D. Nguyen, Raj Vaswani, John Zahorjan, Using runtime measured workload characteristics in parallel processor scheduling, in: *Workshop on Job Scheduling Strategies for Parallel Processing*, Springer, 1996, pp. 155–174.
- [5] Jingren Zhou, John Cieslewicz, Kenneth A. Ross, Mihir Shah, Improving database performance on simultaneous multithreading processors, in: Proceedings of the 31st Very Large Data Bases Conference, VLDB, 2005.
- [6] Christina Delimitrou, Christos Kozyrakis, Quasar: Resource-efficient and qos-aware cluster management, *ACM SIGPLAN Not.* 49 (4) (2014) 127–144.
- [7] Konstantin Shvachko, Hairong Kuang, Sanjay Radia, Robert Chansler, The hadoop distributed file system, in: 2010 IEEE 26th Symposium on Mass Storage Systems and Technologies, MSST, 2010, pp. 1–10.
- [8] Matei Zaharia, Mosharaf Chowdhury, Tathagata Das, Ankur Dave, Justin Ma, Murphy McCauly, Michael J. Franklin, Scott Shenker, Ion Stoica, Resilient distributed datasets: A fault-tolerant abstraction for in-memory cluster computing, in: 9th USENIX Symposium on Networked Systems Design and Implementation, NSDI 12, 2012, pp. 15–28.
- [9] S. Lin, M. Paolieri, C. Chou, L. Golubchik, A model-based approach to streamlining distributed training for asynchronous SGD, in: *MASCOTS*, IEEE, 2018.
- [10] Benjamin Hindman, Andy Konwinski, Matei Zaharia, Ali Ghodsi, Anthony D. Joseph, Randy H. Katz, Scott Shenker, Ion Stoica, Mesos: A platform for fine-grained resource sharing in the data center, in: NSDI, vol. 11, 2011, p. 22.
- [11] A. Verma, L. Pedrosa, M. Korupolu, D. Oppenheimer, E. Tune, J. Wilkes, Large-scale cluster management at Google with Borg, in: *EUROSYS*, ACM, 2015.
- [12] Malte Schwarzkopf, Andy Konwinski, Michael Abd-El-Malek, John Wilkes, Omega: flexible, scalable schedulers for large compute clusters, in: Proceedings of the 8th ACM European Conference on Computer Systems, 2013, pp. 351–364.
- [13] Viktor Leis, Peter Boncz, Alfons Kemper, Thomas Neumann, Morsel-driven parallelism: A NUMA-aware query evaluation framework for the many-core age, in: Proceedings of the 2014 ACM SIGMOD International Conference on Management of Data, SIGMOD '14, 2014, pp. 743–754.
- [14] Kunal Agrawal, Jing Li, Kefu Lu, Benjamin Moseley, Scheduling parallel DAG jobs online to minimize average flow time, in: *SIAM Symposium on Discrete Algorithms*, SIAM, 2016, pp. 176–189.
- [15] J. Edmonds, Scheduling in the dark, *Theoret. Comput. Sci.* (1999).
- [16] Stefano Leonardi, Danny Raz, Approximating total flow time on parallel machines, *J. Comput. System Sci.* 73 (6) (2007) 875–891.
- [17] Sungjin Im, Benjamin Moseley, Kirk Pruhs, Eric Torng, Competitively scheduling tasks with intermediate parallelizability, *TOPC* 3 (1) (2016) 4.
- [18] B. Berg, J.P. Dorsman, M. Harchol-Balter, Towards optimality in parallel scheduling, *ACM POMACS* 1 (2) (2018).
- [19] Jeff Edmonds, Donald D. Chinn, Tim Brecht, Xiaotie Deng, Non-clairvoyant multiprocessor scheduling of jobs with changing execution characteristics, *J. Sched.* 6 (3) (2003) 231–250.
- [20] J. Edmonds, K. Pruhs, Scalably scheduling processes with arbitrary speedup curves, in: *SODA '09*, ACM, 2009, pp. 685–692.
- [21] Jeff Edmonds, Sungjin Im, Benjamin Moseley, Online scalable scheduling for the l_k -norms of flow time without conservation of work, in: Proceedings of the Twenty-Second Annual ACM-SIAM Symposium on Discrete Algorithms, SIAM, 2011, pp. 109–119.
- [22] Benjamin Berg, Mor Harchol-Balter, Benjamin Moseley, Weina Wang, Justin Whitehouse, Optimal resource allocation for elastic and inelastic jobs, in: Proceedings of the 32nd ACM Symposium on Parallelism in Algorithms and Architectures, 2020, pp. 75–87.
- [23] SchedMD, SLURM workload manager, 2021, https://slurm.schedmd.com/heterogeneous_jobs.html.
- [24] Jeffrey Dean, Sanjay Ghemawat, MapReduce: simplified data processing on large clusters, *Commun. ACM* 51 (1) (2008) 107–113.
- [25] Stavros Harizopoulos, Vladislav Shkapenyuk, Anastassia Ailamaki, Qpipe: A simultaneously pipelined relational query engine, in: Proceedings of the 2005 ACM SIGMOD International Conference on Management of Data, 2005, pp. 383–394.
- [26] Philipp Moritz, Robert Nishihara, Stephanie Wang, Alexey Tumanov, Richard Liaw, Eric Liang, Melih Elilbol, Zongheng Yang, William Paul, Michael I. Jordan, et al., Ray: A distributed framework for emerging AI applications, in: 13th USENIX Symposium on Operating Systems Design and Implementation, OSDI 18, 2018, pp. 561–577.
- [27] Robert D. Blumofe, Charles E. Leiserson, Space-efficient scheduling of multithreaded computations, *SIAM J. Comput.* 27 (1) (1998) 202–229.
- [28] Robert D. Blumofe, Charles E. Leiserson, Scheduling multithreaded computations by work stealing, *J. ACM* 46 (5) (1999) 720–748.
- [29] Guy E. Blelloch, Phillip B. Gibbons, Yossi Matias, Provably efficient scheduling for languages with fine-grained parallelism, *J. ACM* 46 (2) (1999) 281–321.
- [30] John L. Hennessy, David A. Patterson, *Computer Architecture: A Quantitative Approach*, Elsevier, 2011.
- [31] S. Srinivasan, S. Krishnamoorthy, P. Sadayappan, A robust scheduling strategy for moldable scheduling of parallel jobs, in: Proceedings of the IEEE International Conference on Cluster Computing, CLUSTER '03, 2003, pp. 92–99.
- [32] Gerald Sabin, Matthew Lang, P. Sadayappan, Moldable parallel job scheduling using job efficiency: An iterative approach, in: *Workshop on Job Scheduling Strategies for Parallel Processing*, Springer, 2006, pp. 94–114.
- [33] Nikos Hardavell, Ippokratis Pandis, Ryan Johnson, Naju Mancheril, Anastassia Ailamaki, Babak Falsafi, Database servers on chip multiprocessors: Limitations and opportunities, in: Proceedings of the Biennial Conference on Innovative Data Systems Research, 2007.

- [34] Rares Vernica, Michael J. Carey, Chen Li, Efficient parallel set-similarity joins using mapreduce, in: Proceedings of the 2010 ACM SIGMOD International Conference on Management of Data, 2010, pp. 495–506.
- [35] [Chen He, Ying Lu, David Swanson, Matchmaking: A new mapreduce scheduling technique](#), in: 2011 IEEE Third International Conference on Cloud Computing Technology and Science, IEEE, 2011, pp. 40–47.
- [36] Stavros Harizopoulos, Anastassia Ailamaki, A case for staged database systems, in: CIDR, 2003.
- [37] [Samuli Aalto, Urtzi Ayesta, Rhonda Righter, On the Gittins index in the M/G/1 queue](#), Queueing Syst. 63 (1) (2009) 437–458.
- [38] Donald R. Smith, A new proof of the optimality of the shortest remaining processing time discipline, Oper. Res. 26 (1) (1978) 197–199.
- [39] [Benjamin Berg, Rein Vesilo, Mor Harchol-Balter, heSRPT: Parallel Scheduling to minimize mean slowdown](#), Performance Evaluation 144 (2020) 102–147.