

Acceptance of Waiting Times in High Performance Computing

Stephan Schlagkamp¹ and Johanna Renker²

¹ Robotics Research Institute, TU Dortmund University, Dortmund, Germany
`stephan.schlagkamp@udo.edu`

² Leibniz Research Centre for Working Environment and Human Factors,
Dortmund, Germany
`renker@ifado.de`

Abstract. In high performance computing, users submit computing jobs to a parallel computing infrastructure. Resources for jobs must be allocated according to the job's requirements. At times of high workload, waiting times are unavoidable. Hence, we investigate user acceptance of waiting times in high performance computing. We analyze data provided by Questionnaire for User Habits of Compute Clusters (QUHCC) among high performance computing users ($n = 24$). On the one hand, the results indicate that increasing processing times of jobs lead to greater acceptance of waiting times. On the other hand, the percental difference between processing times and waiting times decreases for increasing job lengths. We suggest that scheduling strategies should respect our findings.

Keywords: User satisfaction · Service levels · High performance computing · Waiting times

1 Introduction

In high performance computing (HPC), different users submit jobs to a parallel computing infrastructure. A job j is characterized by many different properties. This work considers its length r_j (the running time on the HPC environment) and its waiting time w_j . Waiting time incurs in a situation of high workload, when users submit more work than the system can process at a time. In such situation, the HPC system needs strategies to decide which job among all submitted jobs will be processed next. Literature discusses different optimization goals for HPC environments, e.g., cost optimality according to energy consumption [1]. Since human users submit jobs to a certain HPC environment, we consider their satisfaction with waiting times in detail. Therefore, we focus on user satisfaction in high performance computing according to waiting times. Two different questions occur: what is a good criterion to optimize user satisfaction according to waiting times and how does this criterion influence user satisfaction. Hence, we want to find an *acceptable* way to schedule jobs in situations of high workload.

The structure of this work is as follows. The methodology of this study is in Sect. 2. Furthermore, Sect. 3 presents results a conducted survey, where we also

argue for slowdown as an appropriate metric. In Sect. 4, we discuss the correlation between user satisfaction and slowdown. This work ends with a brief conclusion in Sect. 5.

2 Methodology

In this study, we analyze the answers provided by 24 participants, who answered QUHCC [2]. Beside other methods, Pruyn and Smith use such methodology in their research on customer's reaction to waiting times as well [3]. We perform statistical analyses in form of boxplots and empirical cumulative distribution functions (CDF) to argue for an increasing acceptance of waiting times according to job lengths. We show that acceptances of waiting times in HPC environments are exponentially distributed. Therefore, we analyze the *slowdown* of jobs and rate the quality of different slowdown values.

3 Acceptance of Waiting Times

Questionnaire QUHCC contains six questions asking participants for their acceptance of waiting times for jobs of different lengths (cf. scale Acceptance of Waiting Times (AWT)). We analyze the answers given in a survey among different HPC users and call them *observations*.

Definition 1 (Observation). *An observation $o = (p_j, w_j)$ represents the accepted waiting time w_j for a job length p_j . It is based on an answer provided by a participant.*

We cluster the observations according to the different questions into six sets and sort them according to the assumed job lengths:

- INT: interactive jobs
- 10M: ten minutes
- SHO: short jobs (up to four hours)
- 180M: 180 min.
- MED: medium (one to three days)
- LAR: large (longer than three days)

For sets SHO and MED, we chose the average job length of 120 min and two days, which is equal to 2,880 min, minutes as reference. Figure 1 visualizes all observations as boxplots, Table 1 presents means and medians. The figure reveals an increasing acceptance of waiting times for increasing job lengths. The medians increase from 12.5 min to 10,080 min, and the means increase monotonically from 61.2 min to 9,516.5 min.

Based on this observation, we conclude that growing job lengths cause a greater acceptance of delayed results, which we can exploit in terms of user satisfying scheduling in HPC.

We continue the analysis by introducing slowdown, e.g., [4]. Slowdown describes the relative factor between an optimal response time with zero waiting time and the actual response time as the sum of processing and waiting time.

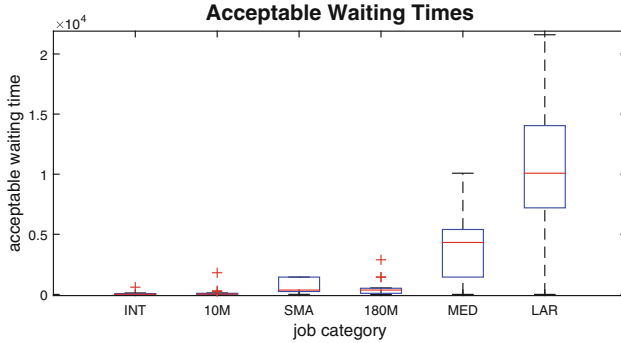


Fig. 1. Acceptable waiting times for six different sets.

Table 1. Mean and median values of acceptable waiting times for six different sets.

	INT	10M	SHO	180M	MED	LAR
median	12.5	30.0	360.0	360.0	4320.0	10,080.0
mean	61.2	171.9	636.7	553.8	3,863.5	9,516.5

Definition 2 (Slowdown). *The slowdown of job j with processing time $p_j > 0$ and waiting time $w_j \geq 0$ is $SD(j) = \frac{p_j + w_j}{p_j}$.*

Therefore, slowdown allows rating the quality of a schedule independent of specific jobs’ lengths and waiting times. Figure 2 depicts the slowdowns for observation sets 10M, SHO, 180M, and LAR as boxplots. Table 2 presents means and medians of slowdowns of each set. We do not consider INT and LAR because of the unspecified exact job lengths disallowing calculation of slowdown. While the acceptance of waiting times increases for larger jobs, the according slowdown decreases. Due to the comparability of observation sets using slowdown, we use this metric to rate the quality of service in the next section.

Table 2. Mean and median values of observations in 10M, SMA, MED, and 180M.

	10M	SMA	180M	MED
mean	18.1938	6.3062	4.0764	2.3415
median	4.0000	4.0000	3.0000	2.5000

4 Quality of Service Regarding Slowdown

By means of *acceptability of slowdown* $A(s, O)$, we calculate the percentage share of observations applying to that slowdown.

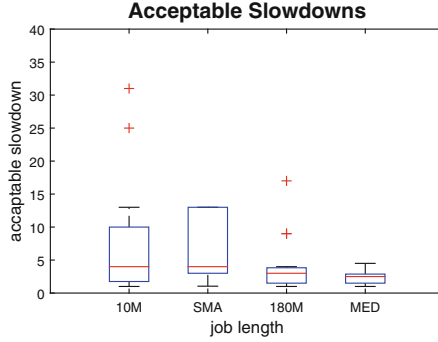


Fig. 2. Acceptable slowdowns for categories 10M, SMA, MED, and 180M.

Definition 3 (Acceptability of Slowdown). *Given a set of observations $O = \{(p_j, w_j)\}$, the acceptability of slowdown s is the amount of observations meeting the slowdown, i.e.,*

$$A(s, O) = \frac{|\{SD(o) \leq s \mid o \in O\}|}{|O|}. \quad (1)$$

Figure 3 depicts the acceptability of slowdowns for all observations as an empirical CDF, as well as an exponential CDF fit, according to an exponential probability density function

$$f(x|\mu) = \frac{1}{\mu} e^{\frac{-x}{\mu}}. \quad (2)$$

To find a suitable fit, we shift the slowdowns by minus one. This calculation is necessary because a minimal value of slowdown is one, while it is zero for a distribution function. Applying least squares, we obtain a parameter of $\mu = 3.85$. We then shift the exponential CDF by one again (cf. Fig. 3). Note that slowdowns greater than 20 are adjusted to 20, since these observations strongly influence the fitting and are not suitable as an optimization objective in practical application.

Beside that, Table 3 presents the acceptability of slowdowns for all sets ranging from 1.0 to 4.0. Set ALL is the union of all sets considered. Based on these plots and analyses, we discuss two different forms of service-levels for HPC environments.

4.1 Single Service Level

To treat all users and jobs equally, we pay special attention to the data of set ALL. Defining a single service quality, we suggest to allow a slowdown between 1.5 and 2.0. For ALL, the amount of satisfied users lays between 74.4% and 85.9%, meaning that almost 3/4 of acceptability is met for 2.0 and less than

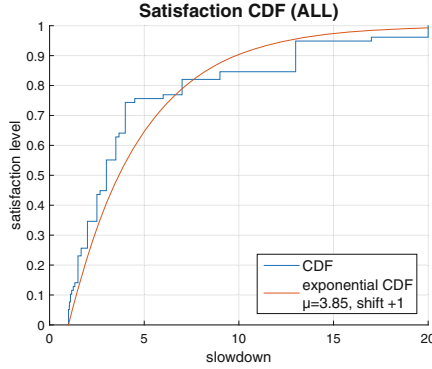


Fig. 3. Acceptability of slowdowns in observations from 10M, SMA, MED, and 180M combined.

Table 3. Satisfaction-Levels in 10M, SMA, MED, and 180M according to different slowdowns.

	1.0	1.5	2.0	2.5	3.0	3.5	4.0
ALL	100 %	85.9 %	74.4 %	65.4 %	55.1 %	44.9 %	35.9 %
10M	100 %	87.5 %	75.0 %	62.5 %	62.5 %	62.5 %	62.5 %
SMA	100 %	95.7 %	87.0 %	78.3 %	78.3 %	69.6 %	52.2 %
180M	100 %	75.0 %	62.5 %	62.5 %	56.3 %	31.3 %	25.0 %
MED	100 %	82.6 %	69.6 %	56.5 %	26.1 %	17.4 %	8.7 %

1/8 of acceptability are not met for a slowdown of 1.5. Nearly the same holds for the fitted exponential CDF with acceptance values between 72.2 % and 85.0 %. Nevertheless, considering set 180M only gains acceptability between 62.5 % to 75.0 %, which is not in the demanded interval. Therefore, we consider multiple service levels next.

4.2 Multiple Service Levels

Beside a single criterion, a scheduling strategy could also pay respect to different job lengths. As analyzed in Sect. 3, the accepted slowdown decreases for increasing job lengths. Therefore, we suggest usage of function (2) to rate acceptability represents a monotonically increasing way to penalize slowdown.

5 Conclusion and Future Work

In this study, we analyzed the data provided by QUHCC asking users for their acceptance of waiting times. The data indicates a growing acceptance of waiting times according to jobs' length, although the relative acceptable waiting time

decreases. Regarding slowdown as a criterion to generalize the accepted waiting times, we argued for a slowdown factor of 1.5–2.0 to suit a tolerable level of users. Additionally, we fitted an exponential distribution to the CDF of slowdowns.

Future work should analyze further influences on user satisfaction, e.g., this work does not respect the point in time when a job ends. Jobs finishing on weekends or in the middle of the night may not suite the users working habits. Such jobs should have finished earlier or a scheduler may exploit, that results are needed on the following monday.

Since the data was collected in a survey among HPC users at TU Dortmund University ($n = 24$), further surveys with more participants at other universities or institutes could support our findings.

References

1. Tang, Q., Gupta, S.K.S., Varsamopoulos, G.: Energy-efficient thermal-aware task scheduling for homogeneous high-performance computing data centers: a cyber-physical approach. *IEEE Trans. Parallel Distrib. Syst.* **19**(11), 1458–1472 (2008)
2. Renker, J., Schlagkamp, S.: QUHCC: questionnaire for user habits of compute clusters. In: *HCI International 2015 - Posters' Extended Abstracts* (in press)
3. Pruyn, A., Smidts, A.: Effects of waiting on the satisfaction with the service: beyond objective time measures. *Int. J. Res. Mark.* **15**(4), 321–334 (1998)
4. Feitelson, D.G.: Metrics for parallel job scheduling and their convergence. In: Feitelson, D.G., Rudolph, L. (eds.) *JSSPP 2001. LNCS*, vol. 2221, p. 188. Springer, Heidelberg (2001)