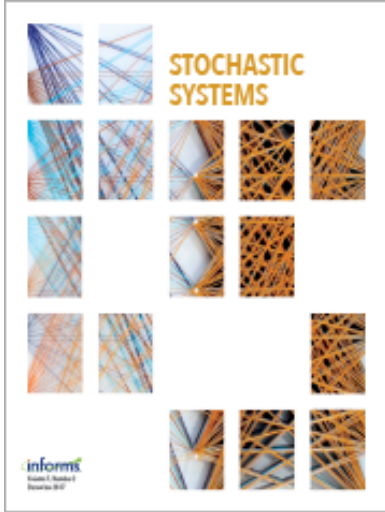




Stochastic Systems

Publication details, including instructions for authors and subscription information:
<http://pubsonline.informs.org>

Asymptotic Optimality of Switched Control Policies in a Simple Parallel Server System Under an Extended Heavy Traffic Condition

Rami Atar; , Eyal Castiel; , Martin I. Reiman

To cite this article:

Rami Atar; , Eyal Castiel; , Martin I. Reiman (2024) Asymptotic Optimality of Switched Control Policies in a Simple Parallel Server System Under an Extended Heavy Traffic Condition. *Stochastic Systems* 14(4):403-454. <https://doi.org/10.1287/stsy.2022.0022>

This work is licensed under a Creative Commons Attribution 4.0 International License. You are free to copy, distribute, transmit and adapt this work, but you must attribute this work as “*Stochastic Systems*. Copyright © 2024 The Author(s). <https://doi.org/10.1287/stsy.2022.0022>, used under a Creative Commons Attribution License: <https://creativecommons.org/licenses/by/4.0/>.”

Copyright © 2024 The Author(s)

Please scroll down for article—it is on subsequent pages



With 12,500 members from nearly 90 countries, INFORMS is the largest international association of operations research (O.R.) and analytics professionals and students. INFORMS provides unique networking and learning opportunities for individual professionals, and organizations of all types and sizes, to better understand and use O.R. and analytics tools and methods to transform strategic visions and achieve better outcomes.

For more information on INFORMS, its publications, membership, or meetings visit <http://www.informs.org>

Asymptotic Optimality of Switched Control Policies in a Simple Parallel Server System Under an Extended Heavy Traffic Condition

Rami Atar,^{a,*} Eyal Castiel,^a Martin I. Reiman^b

^aViterbi Faculty of Electrical and Computer Engineering, Technion, Haifa 32000, Israel; ^bDepartment of Industrial Engineering and Operations Research, Columbia University, New York, New York 10027

*Corresponding author

Contact: rami@technion.ac.il,  <https://orcid.org/0000-0001-8341-5049> (RA); eyal.ca@campus.technion.ac.il,  <https://orcid.org/0000-0002-0800-583X> (EC); martyreiman@gmail.com,  <https://orcid.org/0000-0003-4919-2894> (MIR)

Received: July 7, 2022

Revised: November 23, 2023

Accepted: May 7, 2024

Published Online in Articles in Advance:
July 26, 2024

<https://doi.org/10.1287/stsy.2022.0022>

Copyright: © 2024 The Author(s)

Abstract. This paper studies a two-class, two-server parallel server system under the recently introduced extended heavy traffic condition, which states that the underlying “static allocation” linear program (LP) is critical, but does not require that it has a unique solution. The main result is the construction of policies that asymptotically achieve previously proved a lower bound, on an expected discounted linear combination of diffusion-scaled queue lengths and are therefore asymptotically optimal (AO). Each extreme point solution to the LP determines a control mode—that is, a set of activities (class-server pairs) that are operational. When there are multiple solutions, these modes can be selected dynamically. It is shown that the number of modes required for AO is either one or two. In the latter case, there is a switching point in the (normalized) workload domain, characterized in terms of a free boundary problem. Our policies are defined by identifying pairs of elementary policies and switching between them at this switching point. They provide the first example in the heavy traffic literature where weak limits under an AO policy are given by a diffusion process where both the drift and diffusion coefficients are discontinuous.



Open Access Statement: This work is licensed under a Creative Commons Attribution 4.0 International License. You are free to copy, distribute, transmit and adapt this work, but you must attribute this work as “Stochastic Systems. Copyright © 2024 The Author(s). <https://doi.org/10.1287/stsy.2022.0022>, used under a Creative Commons Attribution License: <https://creativecommons.org/licenses/by/4.0/>.”

Funding: R. Atar is supported by the Israel Science Foundation [Grant 1035/20].

Keywords: parallel server systems • decomposable service rates • extended heavy traffic condition • dynamic graph of basic activities • switched control systems • diffusion with discontinuous coefficients

1. Introduction

1.1. Background

Parallel server systems (PSSs) are queueing control problems in which a number of servers offer service to customers of different classes, and choices as to which customer class each server is dedicated to are made dynamically. Since its introduction in Harrison and López (1999), its study in heavy traffic has attracted much attention, due to its simple structure, its practical significance, and the theoretical challenges it poses. The problem formulation in Harrison and López (1999) includes a key assumption, referred to as the *heavy traffic condition* (HTC), which states that an underlying “static allocation” linear program (LP) satisfies a critical load condition and that it has a unique solution. Whereas critical load is universally considered a defining condition of any notion of heavy traffic, uniqueness of solutions has been assumed mainly because it simplifies the mathematical treatment. The *extended HTC* (EHTC), which merely states that the LP is at criticality, but does not require uniqueness, has recently been introduced in Atar et al. (2024) in order to address a considerably broader notion of heavy traffic. The main result of Atar et al. (2024) is a lower bound on the asymptotically achievable cost in a general PSS under the EHTC. This paper focuses on the two-class, two-server PSS referred to in this introduction as the 2×2 PSS, which is the simplest case in which the EHTC is strictly broader than the HTC. The goal is to complement the results of Atar et al. (2024), in this case by constructing policies that asymptotically achieve the lower bound, which, hence, are asymptotically optimal (AO) in heavy traffic.

The structure of the 2×2 PSS is as follows. Each of the two servers is capable of serving each of the two classes. The classes (respectively, servers) are usually indexed using the symbol i (respectively, k) and activities, namely, class-server pairs, by $j = (i, k)$. Arriving customers await service in class-based queues and, upon receiving a single service, leave the system. The control decisions consist of routing (determining which server serves each customer) and sequencing (determining the order in which they are served). The rates of arrivals of customers of the two classes are denoted by λ_i^n , $i = 1, 2$, and the rates of service at each of the four activities are denoted by μ_{ik}^n , where n denotes the usual heavy traffic parameter. These rates are assumed to be asymptotic to $\lambda_i n + \hat{\lambda}_i n^{1/2}$ and $\mu_{ik} n + \hat{\mu}_{ik} n^{1/2}$, for some given $\lambda_i, \hat{\lambda}_i, \mu_{ik}, \hat{\mu}_{ik}$. The cost consists of an expected infinite horizon discounted linear combination of the two queue lengths and is rescaled at the diffusion scale.

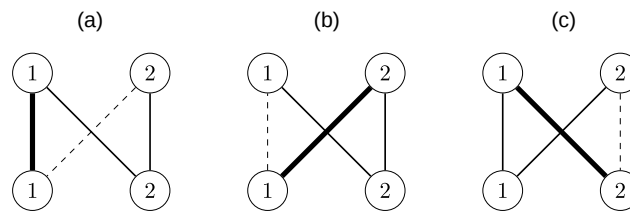
Whereas the cost and, consequently, the notion of AO are set up at the diffusion scale, the underlying LP alluded to above addresses the behavior of the PSS at the fluid, or law-of-large-numbers (LLN), scale. Posed in terms of the first-order parameters, λ_i, μ_{ik} , it is concerned with the mean fraction of time devoted by each server to each class. When the LP has a unique solution, at least one activity is *nonbasic*, in the sense that the fraction allocated to it is zero. The so-called *graph of basic activities* (GBA), formed by the activities with positive allocation fraction, is static. In this case, the critical load condition dictates that any policy not adhering to this solution, in the sense of effort allocation, causes the total queue length to blow up and, in particular, cannot be AO. Under policies that adhere to this solution, the LLN assures that the aforementioned fractions of effort converge to those given by the LP solution (a necessary, but certainly not sufficient, condition for AO). When there are multiple LP solutions, a result from Atar et al. (2024) states that for the 2×2 PSS, the space of solutions, denoted by $\mathcal{S}_{LP}$, forms a line segment $\text{ch}(\xi^{*,1}, \xi^{*,2})$ in the space of 2×2 matrices (where “ch” denotes the convex hull). In each of the two extremal solutions, $\xi^{*,1}, \xi^{*,2}$ there is again at least one nonbasic activity. We refer to these two extreme points as *control modes*, or simply *modes*. For similar reasons, any policy that does not lead to an unbounded cost should keep the system critically loaded at all times, and, thus, asymptotically, the fractions of effort will vary dynamically within $\mathcal{S}_{LP}$. In the control literature, a control process that takes values only at the vertices of the action space is called a bang-bang control. The analogue of this notion in our setting is a policy for which the limiting fractions of effort take values only in $\{\xi^{*,1}, \xi^{*,2}\}$, switching between the two extremal solutions. Some of the policies introduced in this paper are designed to act that way.

Contrary to the setting where the HTC holds, it is impossible to construct an AO policy based only on the first-order data under the EHTC. A second-order approximation of the PSS, which is often referred to as a *Brownian control problem* (BCP), is required. The BCP represents a diffusion limit of the PSS, in which Brownian motion (BM) replaces stochastic fluctuations associated with cumulative arrival and service processes. Closely related to the BCP is another diffusion control problem, called a *workload control problem* (WCP). Obtained by a certain projection of the BCP, it is a control problem in which the process is one-dimensional, representing the total workload asymptotics. The structure of the WCP obtained is quite simple to describe. The state process is a reflected diffusion on $\mathbb{R}_+$, with controlled drift and diffusion coefficients, $b = b(\xi), \sigma = \sigma(\xi)$, where the control process, $\xi = \xi_t$, takes values in $\mathcal{S}_{LP}$, and $\xi \mapsto (b(\xi), \sigma(\xi)^2)$ is an affine map. The cost is given as an expected discounted version of the state process itself. By a standard argument based on the Hamilton-Jacobi-Bellman (HJB) equation, there exists an optimal bang-bang control for the WCP. There can therefore be two possibilities for the WCP solution: the single-mode case, where one of the modes is always used, and the dual-mode case, where both modes are used by the optimal control in different parts of the state space. Note that in this case, the GBA can be changed dynamically. The HJB equation also reveals the structure of the feedback function from state to control. This particular HJB equation was solved in Sheng (1978). It was shown that in the dual-mode case, there is a switching point $z^* \in (0, \infty)$, such that one of the modes is used when the state is below z^* and the other otherwise. The HJB equation can be viewed, in this case, as an equation involving a free boundary, in which the solution is a pair, where one component is the value function and the other is z^* . The results of Sheng (1978) also characterize z^* as the unique solution to an explicit equation, as well as a solution to the HJB equation.

Our policies are obtained by “translating” the WCP solution. In the case of a single mode, the prescribed policy corresponds either to a threshold policy of the form that first appeared in Bell and Williams (2001) (see below) or a simple priority policy, depending on the mode used and the cost. In the dual-mode case, pairs of elementary policies are identified, which are combined together to form switched control policies, so that one is active when the normalized workload process is below the switching point and the other above it. In each case, the policy is designed to meet the target allocation efforts determined by the corresponding mode, and the set of operational activities is restricted by the corresponding GBA.

The paper closest to ours is the aforementioned Bell and Williams (2001), which studies a two-server, two-class PSS with three activities. This PSS is known as an “N” network because upon relabeling, the activities are given by (1, 1), (1, 2), and (2, 2), forming the symbol N. (See Figure 1(a).) In this network, the number of solutions to the

Figure 1. The Full and Nonbasic Activities Are Shown in Thick and Dashed Lines, Respectively



Notes. Graph (a) corresponds to a mode in canonical form. Graphs (a) and (b) correspond to a pair of modes where the nonbasic activity switches a class, whereas in graphs (a) and (c), it switches a server. The pair (b) and (c), in which the nonbasic activity switches both a class and a server, is neither class- nor server-switched.

LP cannot exceed one, and, thus, the requirement of a unique solution does not pose a restriction. In an earlier work, Harrison (1998), it had been observed that when the larger “ $c\mu$ ” value is in class 1, the BCP solution suggests that the queue length of class 1 customers and the idleness process at server 1 should both converge to zero at the diffusion scale and that a simple priority policy does not achieve this. In Bell and Williams (2001), this was addressed by putting a threshold on class 1 queue length that, when exceeded, server 2 prioritizes class 1, and otherwise it prioritizes class 2. The size of the threshold must converge to zero at the diffusion scale so as to achieve the first goal. To achieve AO of a threshold policy with logarithmic (in n) size threshold, as used in Bell and Williams (2001), the interarrival and service times are assumed there to possess exponential moments. (More on the history of the problem and the works that contributed to its development can be read in Atar et al. 2024.)

As already mentioned, one of the policies we implement is a threshold policy similar to the one used by Bell and Williams (2001). However, our assumptions are positioned differently with respect to the threshold-moment tradeoff, assuming a larger (still $o(\sqrt{n})$), polynomial size threshold, but requiring only a polynomial moment assumption. We assume $2 + \varepsilon$ moment assumptions for all of our policies except the single-mode threshold policy and the dual-mode policies that employ the threshold policy when the workload is above z^* . For these, a finite $\mathbf{m}_0$ -th moment is assumed, where the number $4 < \mathbf{m}_0 < 5$ is indicated explicitly. Another difference between our results and those of Bell and Williams (2001) is that our policies do not use preemption. Although the policy introduced in Bell and Williams (2001) uses preemption, it is plausible that an analogous nonpreemptive policy is also AO under similar conditions; see Corollary 2.15 and Remark 2.16. In this paper, our choice not to use preemption leads to nontrivial issues in the dual-mode case. Instead of a simple switching between elementary policies when the workload crosses z^* , it is sometimes the case that one must wait for a particular server to become available before switching. This is described in Section 5.1.3.

It is also worth mentioning that we have argued in Atar et al. (2024) that the AO of the threshold policy from Bell and Williams (2001) extends beyond the HTC to the case of multiple solutions and a single mode (under some assumptions, which include the existence of exponential moments).

Besides the objective to break the uniqueness barrier, an additional source of motivation for this work stems from the relation between nonuniqueness and service rate decomposability. As stated in Lemma 2.1, for the 2×2 PSS, the LP exhibits multiple solutions if and only if the service rates decompose as $\mu_{ik} = \alpha_i \beta_k$. Service rates decompose this way when the mean size of a job is characteristic to the class (and then α_i is the reciprocal mean) and each server has its own processing speed (here given by β_k). As the HTC does not hold under decomposability, this important class of service rates has been left out by earlier work.

1.2. Results

The description of the policies given above is only a sketch. There are nontrivial issues that arise regarding the need to “patch” two policy types, requiring us to slightly modify the policies, where the details differ from one pairing to another.

The main result states that, under the prescribed policies, the rescaled workload process converges in law to the diffusion process that solves the WCP, and these policies are AO. As far as convergence is concerned, in addition to the “standard” issues involved in proving state-space collapse, we need to deal with issues related to switching control modes at z^* . Moreover, to obtain AO from weak convergence, uniform integrability needs to be established, and it is here where the $2 + \varepsilon$ and $\mathbf{m}_0$ moment assumptions are used.

An approach to proving convergence to a diffusion with discontinuous coefficients, addressing especially the technicalities involved with the discontinuity of the diffusion coefficient, was developed in Krylov and Liptser (2002), going beyond the general framework for convergence of semimartingales, such as that from Liptser and

Shiryaev (1977). Whereas the tools from Krylov and Liptser (2002) are not directly applicable in our setting, an argument that, as in Krylov and Liptser (2002), shows that the time spent near the discontinuity set is negligible is also at the basis of our proof. The paper Krylov and Liptser (2002) also gives an example of a queueing model whose scaling limit yields a diffusion process with discontinuities in both drift and diffusion coefficients. Our dual-mode case provides what seems to be the first example where this occurs under an AO policy of a queueing control problem (for AO in heavy traffic leading to discontinuity in the drift only, see Atar and Lev-Ari 2018).

1.3. Organization of the Paper

In Section 2.1, we describe our model and the control problem associated with it in more detail. The LP and the extended heavy traffic condition are introduced in Section 2.2, and preliminary results about the LP from Atar et al. (2024) are stated. In Section 2.3, the WCP and the associated HJB equation are introduced, and in Proposition 2.6, it is stated that there exists a unique classical solution to the HJB equation. This proposition also provides a condition that determines whether an optimal solution to the WCP must employ a single mode or two modes (not to be confused with the number of modes in the space of LP solutions, which is always two under multiplicity) and asserts that in the dual-mode case, there exists a single switching point z^* in workload space. We also present in this section the lower bound from Atar et al. (2024) stated in Theorem 2.4. The main result is stated in Section 2.4. The definitions of the proposed policies appear first, and then, in Theorem 2.13, the weak convergence and AO results are stated. Numerical results are presented in Section 2.5.

In Section 3, we state and prove some results related to the static allocation LP, providing, in particular, explicit expressions for the extreme points of the set of optimal solutions. Development of the WCP is carried out in Section 4. Preliminary results proved in Atar et al. (2024) in a general case are included in this section. This section also contains proofs of results related to the HJB equation, some of which rely on Sheng (1978).

The proof of our main result, Theorem 2.13, is the subject of Section 5. In Section 5.1, we present the general scheme for proving Theorem 2.13; the weak convergence result is stated in Theorem 5.1. We then present four propositions that are used for the proof. Each proposition corresponds to a specific section and step of the proof. Proposition 5.3, in Section 5.2.2, proves uniform integrability; state space collapse is proved in Proposition 5.4 in Section 5.2.3; a key nonidling property is proved in Proposition 5.5 in Section 5.2.4; and a “fast switching” property, showing the aforementioned property that the process spends asymptotically negligible time near the discontinuity, is proved in Proposition 5.6 in Section 5.2.6. Appendix A contains proofs of several lemmas stated earlier. Appendix B presents expressions from Sheng (1981) for the solution to the HJB equation. Finally, Appendix C contains a result providing symmetry conditions under which a single mode control is optimal.

1.4. Notation

$\mathbb{N}$, $\mathbb{R}$, and $\mathbb{R}_+$ are the sets of natural, real, and, respectively, nonnegative real numbers. For $a, b \in \mathbb{R}$, $a \vee b$ and $a \wedge b$ denote the maximum and minimum of a and b , respectively, and $a^+ = a \vee 0$. For a set A , $\mathbb{1}_A$ denotes its indicator function. For $f: \mathbb{R}_+ \rightarrow \mathbb{R}$, and $t, \delta > 0$, $\|f\|_t = \sup_{s \in [0, t]} |f(s)|$ and

$$w_t(f, \delta) = \sup\{|f(s_1) - f(s_2)| : 0 \leq s_1 \leq s_2 \leq (s_1 + \delta) \wedge t\}.$$

For $0 \leq s \leq t$, the notation $f[s, t]$ stands for $f(t) - f(s)$. For real-valued functions and processes, the notation $X(t)$ is used interchangeably with X_t . Given a Polish space E , $C_E[0, \infty)$ and $D_E[0, \infty)$ denote the spaces of E -valued, continuous and, respectively, càdlàg functions on $[0, \infty)$, equipped with the topology of convergence uniformly on compacts and, respectively, the J_1 topology. Denote by $C_{\mathbb{R}}^+[0, \infty)$ and $D_{\mathbb{R}}^+[0, \infty)$ the subset of $C_{\mathbb{R}}[0, \infty)$ and, respectively, $D_{\mathbb{R}}[0, \infty)$, of nonnegative, nondecreasing functions and by $C_{\mathbb{R}}^{0,+}[0, \infty)$ the subset of $C_{\mathbb{R}}^+[0, \infty)$ of functions that are null at zero. Write $X_n \Rightarrow X$ for convergence in law. A sequence of processes with sample paths in $D_E[0, \infty)$ is said to be C -tight if it is tight and the limit of every weakly convergent subsequence has sample paths in $C_E[0, \infty)$ almost surely (a.s.). The letter c denotes a deterministic constant whose value may change from one appearance to another.

2. Model and Main Results

The setup and results presented in this section have quite a few ingredients, as already mentioned in the introduction: the LP and modes, the diffusion scaling, the WCP and HJB equation, the switching point z^* , and threshold and switching policies. Before going into the details, we provide a roadmap. The two-class, two-server system, introduced in Section 2.1, is indexed by n and is observed at the diffusion scale. The assignment of jobs from the different classes to different servers is regarded as a control policy. A cost functional is considered, given by the expected discounted weighted sum of queue lengths, also rescaled at the diffusion scale; see Section 2.6. The weights of the queue lengths are given by a vector (h_1, h_2) .

In order to formulate a policy-independent condition for heavy traffic, an LP (2.7) is introduced, involving first-order arrival and service rates, where the variable to be determined is a 2×2 allocation matrix describing the (first-order) fraction of effort each server dedicates to each class. It is assumed to be at criticality, in the sense that servers are fully occupied, but queue lengths stay balanced. In that regard, we adopt the heavy traffic notion of Harrison and López (1999) and many other papers that followed, with one important distinction: We do not assume that there is a unique allocation matrix that maintains criticality—that is, there may be multiple LP solutions. When there are multiple solutions, they are all given as convex combinations of two allocation matrices (Lemma 2.1), which we call modes. Because earlier work has addressed the unique solution case, we only treat the case of multiple solutions (Assumption 2.2). A result from Atar et al. (2024) states that under this assumption, the first-order rates μ_{ik} are necessarily decomposable as $\alpha_i \beta_k$. Each of the modes induces a so-called graph of basic activities, as in Figure 1. The parameters α_i and the structure of the graphs influence the type of control policy to be proposed.

The WCP (2.16) and (2.17) is a control problem for a diffusion process in dimension 1, in which both the drift and the diffusion coefficient are controlled by a control process that takes values in the space of allocation matrices. The values of these coefficients, evaluated at the two modes, are denoted b_m and, respectively, σ_m , $m = 1, 2$. The WCP formally describes the control problem associated with the PSS at the diffusion limit, and its control process represents the dynamic selection of the mode at which the PSS operates. The significance of the WCP has two aspects. First, as was shown in Atar et al. (2024), its solution gives a lower bound on the PSS cost asymptotics under any sequence of control policies (Theorem 2.4). Accordingly, any policy that achieves this bound is AO. Second, it suggests how AO policies should be structured. In particular: (a) State space collapse should hold, which, in our setting, involves two properties that should hold up to a level negligible at diffusion scale: (i) Both servers should be busy whenever there is any work in the system, and (ii) all queue length should be kept in the class i that minimizes $h_i \alpha_i$. Roughly speaking, (ii) implies that the class that maximizes $h_i \alpha_i$ should be prioritized. This type of sequencing policy is known as a $c\mu$ rule. (But, as shown by Harrison 1998 and Bell and Williams 2001, something more than a simple priority policy may be needed here in order to accomplish (i).) (b) When, for either $(m, m') = (1, 2)$ or $(m, m') = (2, 1)$, one has $b_m \leq b_{m'}$ and $\sigma_m \leq \sigma_{m'}$, only mode m should be used. Otherwise both modes should be used, by dynamically selecting them at different parts of the state space. The partition of the state space is defined via a switching point z^* : When the diffusion-scaled workload is below z^* , the mode m with $b_m \geq b_{m'}$ should be used; otherwise, m' . Determining this switching point is done by solving an HJB equation, (2.18) and (2.19).

The construction of an AO policy must take into account several considerations: whether to operate in one or two modes and, in the latter case, their ordering in workload space; which class has high priority; and the structure of the graph of basic activities for each of the modes. Different types of policies apply in different cases (Definitions 2.9, 2.10, 2.11, and 2.12). The main result states that these policies are AO and that the normalized workload converges weakly to the diffusion process given by the WCP state process under an optimal control (Theorem 2.13).

2.1. Queueing Model, Scaling, and Queueing Control Problem

The model under consideration is as in Atar et al. (2024), specialized to the case of two job classes, two servers, and four activities. We will refer to it as the 2×2 PSS when there is need to distinguish it from the general PSS treated in Atar et al. (2024). The symbol $i \in \{1, 2\}$ is used as a generic index to a class and $k \in \{1, 2\}$ to a server. For a general PSS, an *activity* is a class-server pair (i, k) , where server k is capable of serving class i . In this paper, it is assumed that each server is capable of serving each class; hence, there are four activities. They are labeled by (i, k) or sometimes by $j = (i, k)$.

The model consists of a sequence of systems, indexed by $n \in \mathbb{N}$, that are all defined on one probability space $(\Omega, \mathcal{F}, \mathbb{P})$. For the n th system, one considers the following processes. The processes denoted $A^n = (A_i^n)$ and $S^n = (S_{ik}^n)$ represent arrival and potential service counting processes. That is, $A_i^n(t)$ is the number of arrivals of class i jobs until time t , $i = 1, 2$, and $S_{ik}^n(t)$ is the number of service completions of class i jobs by server k , by the time server k has devoted t units of time to class i , $i = 1, 2$, $k = 1, 2$. Next, $X^n = (X_i^n)$, $I^n = (I_k^n)$, $D^n = (D_{ik}^n)$, and $T^n = (T_{ik}^n)$ denote queue length, cumulative idleness, departure, and cumulative busyness processes. In other words, $X_i^n(t)$ is the number of class i customers in the system at time t , $I_k^n(t)$ is the cumulative time server k has been idle by time t , $D_{ik}^n(t)$ is the number of class i departures from server k , and $T_{ik}^n(t)$ is the cumulative time devoted by server k to class i . The process T_{ik}^n takes the form $T_{ik}(t) = \int_0^t \Xi_{ik}^n(s) ds$, where $\Xi_{ik}^n(t)$ is the fraction of effort devoted by server k to class- i jobs at t . In particular, $\sum_i \Xi_{ik}^n(t) \leq 1$ for every k . Thus, Ξ^n is referred to as the allocation process.

The aforementioned arrival and potential service processes are constructed as follows. Arrival rates λ_i^n and service rates μ_{ik}^n are given, satisfying, for some constants $\lambda_i \in (0, \infty)$, $\mu_{ik} \in (0, \infty)$, $\hat{\lambda}_i \in \mathbb{R}$, $\hat{\mu}_{ik} \in \mathbb{R}$,

$$\begin{aligned}\hat{\lambda}_i^n &:= n^{-1/2}(\lambda_i^n - n\lambda_i) \rightarrow \hat{\lambda}_i, \\ \hat{\mu}_{ik}^n &:= n^{-1/2}(\mu_{ik}^n - n\mu_{ik}) \rightarrow \hat{\mu}_{ik},\end{aligned}$$

as $n \rightarrow \infty$. For each i , a renewal process $\check{A}_i$ is given, with interarrival distribution that has mean one and squared coefficient of variation $0 < C_{\check{A}_i}^2 < \infty$. Similarly, for each (i, k) , a renewal process $\check{S}_{ik}$ is given with mean one interarrival and squared coefficient of variation $0 < C_{\check{S}_{ik}}^2 < \infty$. It is assumed that A^n and S^n are given by

$$A_i^n(t) = \check{A}_i(\lambda_i^n t), \quad S_{ik}^n(t) = \check{S}_{ik}(\mu_{ik}^n t).$$

It is assumed, moreover, that the six processes $\check{A}_i, \check{S}_{ik}$ are mutually independent and have strictly positive interarrival distributions and right-continuous sample paths. The independent and identically distributed interarrivals of $\check{A}_i$ and $\check{S}_{ik}$ are denoted by $\check{a}_i(l)$ and $\check{u}_{ik}(l)$, $l \in \mathbb{N}$, respectively, and those of the accelerated processes A_i^n and S_{ik}^n are given by

$$a_i^n(l) = \frac{1}{\lambda_i^n} \check{a}_i(l), \quad u_{ik}^n(l) = \frac{1}{\mu_{ik}^n} \check{u}_{ik}(l). \quad (2.1)$$

The system is assumed to start empty; that is, $X^n(0) = 0$ for all n . Simple relations between the processes are

$$D_{ik}^n(t) = S_{ik}^n(T_{ik}^n(t)), \quad (2.2)$$

$$X_i^n(t) = A_i^n(t) - \sum_k D_{ik}^n(t), \quad (2.3)$$

$$I_k^n(t) = t - \sum_i T_{ik}^n(t), \quad (2.4)$$

$$\text{the sample paths of } X_i^n \text{ are nonnegative, and those of } I_k^n \text{ are in } C_{\mathbb{R}}^{0,+}[0, \infty). \quad (2.5)$$

The tuple $(\check{A}, \check{S})$ is referred to as the *stochastic primitives*. In our formulation, we will consider T^n as the control process (equivalently, the allocation process Ξ^n may be regarded the control). In view of Equations (2.2), (2.3), and (2.4), given the stochastic primitives, the control uniquely determines the processes D^n, X^n, I^n . Let an additional process be defined on the probability space denoted by $\Upsilon = (\Upsilon(l), l \in \mathbb{N})$, taking values in a Polish space $\mathcal{S}_{\text{rand}}$ and assumed to be independent of the stochastic primitives, for each n (there is no need to let Υ vary with n , as the primitives are all defined on the same probability space). It is included in the model in order to allow the construction of randomized controls; for more details about its potential use, see Atar et al. (2024, remark 2.1.ii).

The process T^n is said to be an *admissible control for the queueing control problem (QCP) for the n -th system* if for each (i, k) , T_{ik}^n has sample paths in $C_{\mathbb{R}}^{0,+}[0, \infty)$ that are 1-Lipschitz, and the associated processes D^n, X^n , and I^n given by (2.2), (2.3), and (2.4) satisfy (2.5); furthermore, T^n is adapted to the filtration $\{\mathcal{F}_t^n\}$ defined by $\mathcal{F}_t^n = \sigma\{(A^n(s), D^n(s), s \in [0, t]), \Upsilon\}$. Denote by $\mathcal{A}^n$ the collection of all admissible controls for the QCP for the n -th system. As argued in Atar et al. (2024, remark 2.1.i), this definition allows for the control to depend on the history of all processes involved in the model (in addition to the auxiliary randomness Υ).

The queue-length process normalized at the diffusion scale is defined by $\hat{X}_i^n(t) = n^{-1/2}X_i^n(t)$. The cost of interest for the n -th system is given by

$$\hat{J}^n(T^n) = \mathbb{E} \int_0^\infty e^{-\gamma t} h(\hat{X}^n(t)) dt, \quad T^n \in \mathcal{A}^n, \quad (2.6)$$

where $\gamma > 0$ and $h(x) = h_1 x_1 + h_2 x_2$, with constants $h_1, h_2 > 0$, and $\hat{X}^n$ is the rescaled queue length process associated with the admissible control T^n . The value for the n -th system is defined by

$$\hat{V}^n = \inf\{\hat{J}^n(T^n) : T^n \in \mathcal{A}^n\}.$$

This completes the description of the queueing models and QCP. The complete set of problem data consists of the stochastic primitives mentioned above and the collection of parameters

$$(\lambda_i), (\mu_{ik}), (\hat{\lambda}_i), (\hat{\mu}_{ik}), (C_{A_i}), (C_{S_{ik}}), \gamma, (h_i).$$

We sometimes refer to $(\lambda_i), (\mu_{ik})$ as the first-order data and to $(\hat{\lambda}_i), (\hat{\mu}_{ik}), (C_{A_i}), (C_{S_{ik}})$ as the second-order data.

2.2. The Linear Program and Extended Heavy Traffic Condition

Given the first-order data $\lambda_i > 0$, $\mu_{ik} > 0$, $i, k = 1, 2$, consider the following linear program for the unknowns $(\xi_{ik}) \in \mathbb{R}_+^{2 \times 2}$ and $\rho \in \mathbb{R}$.

Linear Program. Minimize ρ subject to

$$\begin{cases} \sum_{k=1}^2 \xi_{ik} \mu_{ik} = \lambda_i & i = 1, 2, \\ \sum_{i=1}^2 \xi_{ik} \leq \rho & k = 1, 2, \\ \xi_{ik} \geq 0 & i, k = 1, 2. \end{cases} \quad (2.7)$$

Denote the optimal objective value of (2.7) by ρ^* .

Extended Heavy Traffic Condition. $\rho^* = 1$

The extended heavy traffic condition is broader than the *heavy traffic condition* that has been extensively used in the literature, which requires, in addition to $\rho^* = 1$, that there be a unique corresponding ξ .

Under the EHTC, any solution is of the form $(\xi, 1)$. Let $\mathcal{S}_{\text{LP}}$ denote the subset of $\mathbb{R}_+^{2 \times 2}$, for which the set of all solutions is given by $\mathcal{S}_{\text{LP}} \times \{1\}$. We say that the *EHTC with multiplicity* (EHTCM) holds if the EHTC holds and the LP has multiple solutions (that is, there exist two distinct pairs $(\xi^{(1)}, 1)$ and $(\xi^{(2)}, 1)$ satisfying (2.7)). Following Atar et al. (2024), we say that the service rates μ_{ik} are *decomposable* if $\mu_{ik} = \alpha_i \beta_k$ for all i, k , for some constants α_i and β_k .

A matrix $\xi \in \mathbb{R}_+^{2 \times 2}$ is called *column-stochastic* if $\sum_i \xi_{ik} = 1$ for both $k = 1, 2$. A column-stochastic matrix is called a *mode* if (at least) one of its columns is either $(0, 1)^T$ or $(1, 0)^T$. A mode is said to be *degenerate* if it has more than one zero entry; otherwise, it is said to be *nondegenerate*. A pair of nondegenerate modes is said to be a *class-switched* (server-switched) pair of modes if the zero entries in the two modes are in distinct rows but the same column (respectively, distinct columns but the same row). For example, the graphs in Figure 1, (a) and (b) correspond to a class-switched pair of modes, whereas those in Figure 1, (a) and (c) correspond to a server-switched pair.

The following condition will be referred to as the *nondegeneracy condition*, namely,

$$\lambda_i \neq \mu_{ik} \text{ for all } (i, k) \in \{1, 2\}^2. \quad (2.8)$$

The following is proved in Section 3.

Lemma 2.1. Let the EHTC hold.

1. For any solution $(\xi, 1)$, ξ is column-stochastic.
2. The LP (2.7) has multiple solutions if and only if (μ_{ik}) are decomposable.
3. If the LP has multiple solutions, then there exists a pair of modes $(\xi^{*,1}, \xi^{*,2})$ such that

$$\mathcal{S}_{\text{LP}} = \text{ch}(\{\xi^{*,1}, \xi^{*,2}\}). \quad (2.9)$$

4. If the LP has multiple solutions and the nondegeneracy condition (2.8) holds, then both $\xi^{*,1}$ and $\xi^{*,2}$ of (2.9) are nondegenerate. Moreover, they form either a class-switched or a server-switched pair.

The main result will be proved under the following.

Assumption 2.2.

1. The EHTCM holds.
2. The nondegeneracy condition (2.8) holds.

Note that the case where the EHTC holds, but EHTCM does not hold, is already covered in the work Bell and Williams (2001), although, as mentioned in the introduction, under different assumptions on preemption and moment conditions. (See Corollary 2.15 and Remark 2.16 for implications of our results to the case where uniqueness holds and more on the relation to Bell and Williams 2001 in that case.)

In view of Lemma 2.1(2), under Assumption 2.2, the rates (μ_{ik}) are decomposable. Thus, $\mu_{ik} = \alpha_i \beta_k$, and clearly, there is a degree of freedom in choosing (α_i) and (β_k) . In this paper, we will always assume that they are chosen so that $\sum_k \beta_k = 1$, and it is easy to see that, given (μ_{ik}) , this normalization uniquely determines these parameters.

It is also guaranteed by the lemma that, under Assumption 2.2, the extreme points of $\mathcal{S}_{LP}$ are two nondegenerate modes $\xi^{*,1}, \xi^{*,2}$ forming a class- or a server-switched pair. Once a labeling of these modes has been fixed, we will sometimes slightly abuse the terminology by referring to them as modes 1 and 2, rather than modes $\xi^{*,1}$ and $\xi^{*,2}$.

In earlier work on PSS, under the assumption that the LP has a unique solution $(\xi^*, 1)$, activities are categorized as *basic* or *nonbasic* according to the positivity of the fraction allocated to them by ξ^* ; that is, an activity (i, k) is *basic* if $\xi_{ik}^* > 0$ and *nonbasic* if $\xi_{ik}^* = 0$. We extend this terminology to the case of multiple solutions as follows. For $m \in \{1, 2\}$, an activity (i, k) is said to be *basic in mode m* if the allocation associated to it by this mode does not vanish, namely, $\xi_{ik}^{*,m} > 0$. If $\xi_{ik}^{*,m} = 0$ (respectively, $\xi_{ik}^{*,m} = 1$) it is said to be *nonbasic (respectively, full) in mode m* .

Figure 1 demonstrates the second part of Lemma 2.1(4)—namely, that a pair of modes can be class-switched, as in Figure 1, (a) and (b), or server-switched, as in Figure 1, (a) and (c), but the LP does not give rise to a pair such as Figure 1, (b) and (c). The terms class-switched and server-switched will sometimes be abbreviated as “CS” and “SS.”

A mode is said to be in *canonical form* if its first column is $(1, 0)^T$. It is clear by the definition of a mode that it is always possible to relabel the classes and the servers so that a given mode is in canonical form and that if the mode is nondegenerate, there is only one such relabeling. The graph of a mode in canonical form is shown in Figure 1(a). Because of its resemblance to the symbol N , this form is sometimes called an *N-system*.

If the EHTCM holds, but (2.8) does not—that is, there exist i, k such that $\lambda_i = \mu_{ik}$ —the situation is different: at least one of the modes will be degenerate—that is, have two nonbasic activities (see Lemma 3.1). In the terminology of linear programming, this corresponds to a case where one of the basic solutions of the LP (2.7) is degenerate (cf. Goldfarb and Todd 1989, definition 3.1). The degenerate case is not covered in this paper. The key difficulty in the degenerate case is that, in at least one of the modes, the servers do not communicate in the graph of basic activities, so the pooling required to reach the cost lower bound is not possible. (See Atar et al. 2024, section 2.4 for a more detailed discussion of this issue.)

Under Assumption 2.2, both modes have a single nonbasic activity. Then, given a mode, the graph of basic activities has exactly three edges, and one can speak of the single-activity class (the one associated with only one nonbasic activity), the dual-activity class, and, similarly, the single- and dual-activity server. These terms allow us to refer to the roles of classes and servers in the graph without considering a particular labeling. It is also useful to accompany these terms with matching notation. For a mode ξ , let the single- (respectively, dual-) activity class be denoted by $i_1(\xi)$ (respectively, $i_2(\xi)$), and, similarly, let the single- (respectively, dual-) activity server be denoted by $k_1(\xi)$ (respectively, $k_2(\xi)$).

The following example, which corresponds to examples (A) and (B) in Atar et al. (2024), should help to make the above ideas more concrete.

Example 2.3. Let (μ_{ik}) be given by

$$\mu_{11} = 3, \quad \mu_{12} = 4, \quad \mu_{21} = 6, \quad \mu_{22} = 8.$$

Then, μ_{ik} are decomposable as $\alpha_i \beta_k$, where $\alpha_1 = 7, \alpha_2 = 14, \beta_1 = 3/7$, and $\beta_2 = 4/7$. For (λ_i) , consider two cases—namely,

$$(A) \quad \lambda_1 = 5, \lambda_2 = 4, \quad (B) \quad \lambda_1 = 3.5, \lambda_2 = 7.$$

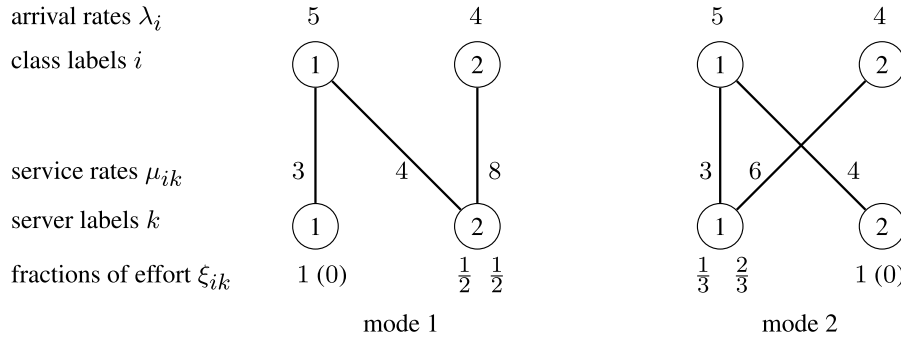
The linear program takes the form: Minimize ρ subject to

$$\begin{cases} 3\xi_{11} + 4\xi_{12} = \lambda_1, \\ 6\xi_{21} + 8\xi_{22} = \lambda_2, \\ \xi_{11} + \xi_{21} \leq \rho, \\ \xi_{12} + \xi_{22} \leq \rho, \\ \min \xi_{ik} \geq 0. \end{cases} \quad (2.10)$$

The solutions to the LP were calculated in Atar et al. (2024), and it was found that, in both cases, $\rho^* = 1$, and the two modes are given by

$$\xi^{*,1} = \left(1, \frac{1}{2}, 0, \frac{1}{2}\right)^T, \quad \xi^{*,2} = \left(\frac{1}{3}, 1, \frac{2}{3}, 0\right)^T,$$

Figure 2. The Two Modes in Case (A)



in case (A) and

$$\xi^{*,1} = \left(1, \frac{1}{8}, 0, \frac{7}{8}\right)^T, \quad \xi^{*,2} = \left(0, \frac{7}{8}, 1, \frac{1}{8}\right)^T,$$

in case (B). In particular, the system is critically loaded, and there are multiple ways to allocate the effort so as to meet the demand. That is, the EHTCM holds. Figures 2 and 3 depict the graphs of basic activities corresponding to these modes. The examples differ in the way the two modes are paired. In case (A) (respectively (resp.)), (B)), the nonbasic activity switches a server (a class). This distinction is of crucial significance to the way AO policies are designed in this paper.

2.3. Workload Control Problem

The WCP was derived and studied in Atar et al. (2024) under the EHTC. We describe this problem in the special case needed here—namely, under the setting of a 2×2 PSS and assuming that the EHTCM holds. In particular, as mentioned above, the parameters (α_i) , (β_k) are uniquely determined by the problem data. Define the workload process and its scaled version as

$$W^n(t) = \sum_i \frac{X_i^n(t)}{\alpha_i}, \quad \hat{W}^n(t) = \sum_i \frac{\hat{X}_i^n(t)}{\alpha_i}. \quad (2.11)$$

(It follows from Atar et al. 2024, lemma 2.4(2) that (2.11) agrees with the definition of W^n and $\hat{W}^n$ given in Atar et al. 2024). Let the process that appears in the definition of the Cost (2.6) be denoted by

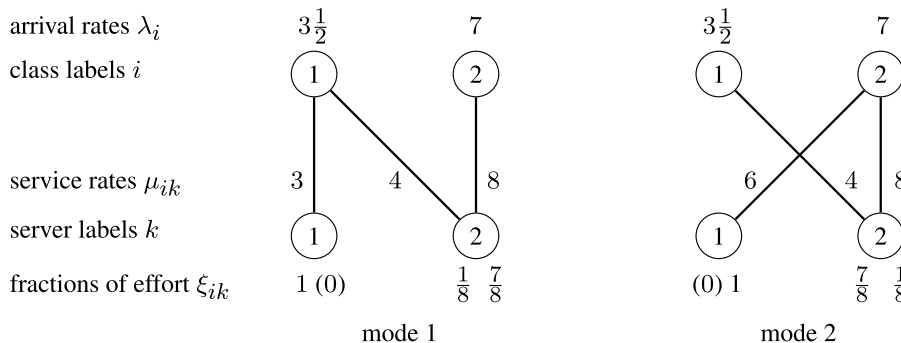
$$\hat{H}_t^n = h(\hat{X}^n(t)) = h_1 \hat{X}_1^n(t) + h_2 \hat{X}_2^n(t). \quad (2.12)$$

Throughout, denote by $p, q \in \{1, 2\}$ the two distinct indices for which

$$h_p \alpha_p \geq h_q \alpha_q, \quad (2.13)$$

where in the special case $h_1 \alpha_1 = h_2 \alpha_2$, set $p = 1$ and $q = 2$. The policies constructed in this paper aim at keeping $\hat{X}_p^n$ close to zero. Hence, we call p the *high-priority class* (HPC) and q the *low-priority class* (LPC). Next, let

Figure 3. The Two Modes in Case (B)



$\sigma_{A,i} = \lambda_i^{1/2} C_{A,i}$, $\sigma_{S,ik} = \mu_{ik}^{1/2} C_{S,ik}$, and

$$b(\xi) = \sum_i \frac{\hat{\lambda}_i - \sum_k \hat{\mu}_{ik} \xi_{ik}}{\alpha_i}, \quad \sigma(\xi)^2 = \sum_i \frac{\sigma_{A,i}^2 + \sum_k \sigma_{S,ik}^2 \xi_{ik}}{\alpha_i^2}, \quad \xi = (\xi_{ik}) \in \mathcal{S}_{LP}. \quad (2.14)$$

Let also

$$b_m = b(\xi^{*,m}), \quad \sigma_m = \sigma(\xi^{*,m}), \quad m = 1, 2. \quad (2.15)$$

It was shown in Atar et al. (2024) that the asymptotics of the pair $(\hat{W}^n, \Xi^n)$ are governed by a state-control pair of processes (Z, Ξ) , where Z is a one-dimensional controlled diffusion given by

$$Z_t = z + \int_0^t b(\Xi_s) ds + \int_0^t \sigma(\Xi_s) dB_s + L_t, \quad (2.16)$$

Ξ is a control process, B is a standard BM (SBM), L is a reflection term at zero, and $z \geq 0$. A precise definition is as follows.

A tuple $\mathfrak{S} = (\Omega', \mathcal{F}', (\mathcal{F}'_t), \mathbb{P}', B, \Xi, Z, L)$ is said to be an *admissible control system for the WCP with initial condition z* if $(\Omega', \mathcal{F}', (\mathcal{F}'_t), \mathbb{P}')$ is a filtered probability space; B , Ξ , Z , and L are processes defined on it; B is a SBM and an $(\mathcal{F}'_t)$ -martingale; Ξ is $(\mathcal{F}'_t)$ -progressively measurable taking values in $\mathbb{R}^{2 \times 2}$ and satisfying $\mathbb{P}'$ (for a.e. t , $\Xi_t \in \mathcal{S}_{LP}$) $= 1$; Z is continuous nonnegative and $(\mathcal{F}'_t)$ -adapted; L has sample paths in $C_{\mathbb{R}^+}^{0,+}[0, \infty)$ and is $(\mathcal{F}'_t)$ -adapted; and Equation (2.16) and the identity $\int_{[0, \infty)} Z_t dL_t = 0$ are satisfied $\mathbb{P}'$ -a.s.

Denoting by $\mathcal{A}_{WCP}(z)$ the collection of all control systems for the WCP with initial condition z , the cost and value are defined as

$$J_{WCP}(z, \mathfrak{S}) = \mathbb{E}_{\mathfrak{S}} \int_0^\infty e^{-\gamma t} Z_t dt, \quad V_{WCP}(z) = \inf\{J_{WCP}(z, \mathfrak{S}) : \mathfrak{S} \in \mathcal{A}_{WCP}(z)\}. \quad (2.17)$$

The significance of the WCP lies in the fact that it provides a lower bound on the large n asymptotics of the n -th system value. More precisely, the main result of Atar et al. (2024), when specialized to the 2×2 PSS, states the following.

Theorem 2.4 (Atar et al. 2024). *Let Assumption 2.2 (multiplicity and nondegeneracy) hold. Then,*

$$\liminf_n \hat{V}^n \geq V_0 := h_q \alpha_q V_{WCP}(0).$$

To prove this result, one only needs to verify that the assumptions from Atar et al. (2024) hold under Assumption 2.2. This is done in Section 3.

Remark 2.5. The general result from Atar et al. (2024) does not assume multiplicity and, moreover, uses different notation based on the formulation of a dual to the LP. In the 2×2 setting with multiplicity, the lower bound from Atar et al. (2024) reduces to the above expression given in terms of the parameters (α_i) in place of the dual.

In view of Theorem 2.4, a sequence $T^n \in \mathcal{A}^n$ of admissible controls for the QCP, also referred to as a *sequence of policies*, is said to be *asymptotically optimal* if

$$\limsup_{n \rightarrow \infty} \hat{J}^n(T^n) \leq V_0.$$

2.3.1. A Heuristic Discussion of the WCP and Sheng's Tortoise–Hare Problem. The WCP is important not only because it provides a lower bound on performance, but also because one can learn from it how to construct a policy for the queueing model that performs near optimality. This problem has been solved completely. Before presenting its solution, we discuss its nature informally. Assume first that the pairs (b_1, σ_1) and (b_2, σ_2) are such that $b_1 < b_2$, while $\sigma_1 = \sigma_2$. Then, it is intuitively clear and can be shown by a simple coupling that J_{WCP} is minimized by always using mode 1, because it has a smaller drift. Next, consider the case where $b_1 = b_2$, but $\sigma_1 < \sigma_2$. Mode 2 has greater variance, and, thus, one expects that the constraining mechanism, causing the diffusion to bounce back from the boundary, will be more active under mode 2 than under mode 1. Hence, again, using mode 1 at all times is optimal. More generally, mode 1 is optimal when $b_1 \leq b_2$ and $\sigma_1 \leq \sigma_2$.

A more interesting case is when $b_1 < b_2$ but $\sigma_1 > \sigma_2$, referred to in Sheng (1978) as the *tortoise–hare problem*. It seems reasonable to use the mode with smaller drift when the diffusion process is far from the origin and switch

to the mode with smaller variance when it is close to the origin, where the aforementioned boundary effect is more prominent.

These heuristic arguments were validated rigorously in Sheng (1978). In particular, in the case $b_1 < b_2$, $\sigma_1 > \sigma_2$, it was shown that there exists a point $z^* \in (0, \infty)$ such that it is optimal to select mode 2 (resp., 1) when Z_t is in $[0, z^*)$ (resp., $[z^*, \infty)$). The identification of z^* and the proof of the result were based on an HJB equation. The precise details are as follows.

2.3.2. The HJB Equation. The value function can be characterized in terms of an HJB equation (see Fleming and Soner 2006 for an introduction to the subject). To present this equation, for $(v_1, v_2, \xi) \in \mathbb{R}^2 \times S_{LP}$, let

$$\bar{\mathbb{H}}(v_1, v_2, \xi) = b(\xi)v_1 + \frac{\sigma(\xi)^2}{2}v_2, \quad \mathbb{H}(v_1, v_2) = \inf_{\xi \in S_{LP}} \bar{\mathbb{H}}(v_1, v_2, \xi) = \min_{\xi \in \{\xi^{*,1}, \xi^{*,2}\}} \bar{\mathbb{H}}(v_1, v_2, \xi),$$

where the identity on the right-hand side (RHS) follows from the fact that both b and σ^2 are affine as a function of ξ , and S_{LP} is the convex hull of $\{\xi^{*,1}, \xi^{*,2}\}$. A classical solution to the HJB equation is a $C^2(\mathbb{R}_+ : \mathbb{R})$ function u satisfying

$$\mathbb{H}(u'(z), u''(z)) + z - \gamma u(z) = 0, \quad z \in (0, \infty), \quad (2.18)$$

and the boundary conditions at 0 and ∞ ,

$$u'(0) = 0, \quad u(z) < c(1 + z), z \in \mathbb{R}_+, \text{ for some constant } c. \quad (2.19)$$

Given a C^2 function u , denote $\mathbb{H}_m^u(z) = \bar{\mathbb{H}}(u'(z), u''(z), \xi^{*,m})$, $m = 1, 2$, and $\mathbb{H}^u(z) = \min_m \mathbb{H}_m^u(z)$.

The following two conditions play an important role in what follows. They correspond to the two cases discussed above and will be referred to as the *single-mode* case and, respectively, the *dual-mode* case (not to be confused with uniqueness and multiplicity of the LP solution):

$$\text{there exist distinct } m, m' \in \{1, 2\} \text{ such that } b_m \leq b_{m'} \text{ and } \sigma_m \leq \sigma_{m'}, \quad (2.20)$$

$$\text{there exist distinct } m, m' \in \{1, 2\} \text{ such that } b_m < b_{m'} \text{ and } \sigma_m > \sigma_{m'}. \quad (2.21)$$

The C^2 smoothness of the value function is tied to the question of existence of a classical solution to the HJB equation. Owing to the uniform ellipticity ($\sigma(\xi)^2 > 0$ at both $\xi = \xi^{*,1}$ and $\xi^{*,2}$), one can show that a classical solution uniquely exists (Atar et al. 2024, proposition 2.5), a type of result that, in the general context of optimal switching of a diffusion process, has been called *the principle of smooth fit*. The results of Sheng (1978) alluded to above shed more light on the specific problem at hand, providing structural properties (parts 2 and 3 of the result below) that are harder to obtain via the general approach.

Proposition 2.6.

1. There exists a unique classical solution u to (2.18) and (2.19). Moreover, $u = V_{WCP}$.
2. If (2.20) holds, then, with m as in (2.20),

$$\mathbb{H}^u(z) = \mathbb{H}_m^u(z) \quad z \in \mathbb{R}_+. \quad (2.22)$$

3. Alternatively, if (2.21) holds, then there exists $z^* \in (0, \infty)$ such that, with the pair (m, m') of (2.21),

$$\mathbb{H}^u(z) = \begin{cases} \mathbb{H}_{m'}^u(z) & z < z^*, \\ \mathbb{H}_m^u(z) & z > z^*. \end{cases} \quad (2.23)$$

The proof of this result appears in Section 4. Some details on the construction from Sheng (1978) (as corrected in Sheng (1981), by which parts 2 and 3 above were proved, appear in Appendix B.

Further terminology is as follows. In the single-mode case, the mode $\xi^{*,m}$ for which m satisfies (2.20) will be referred to as the *active mode* and denoted ξ^A because the above result indicates that it is optimal to always select $\Xi_t = \xi^{*,m}$. In the dual-mode case, the modes $\xi^{*,m'}$ and $\xi^{*,m}$ for which m' and m satisfy (2.21) will be referred to as ξ^L , the *lower-* and, respectively, ξ^H the *higher-workload mode*, and z^* as the *switching point*. These terms refer to the fact that the result suggests that it is optimal to select $\Xi_t = \xi^L$ (respectively, $\Xi_t = \xi^H$) when $Z_t < z^*$ (respectively, $Z_t > z^*$).

Example 2.7. We go back to Example 2.3, focusing on case (A), adding now information on the second-order data. Assume that the squared coefficients of variation $C_{A_i}^2$ and $C_{S_{ik}}^2$ are all one except $C_{S_{11}}^2 = 4$. Set both $\hat{\lambda}_i$ to zero. As for $\hat{\mu}_{ik}$, consider two cases. In case (A1), all $\hat{\mu}_{ik}$ are set to zero. In case (A2), they are all zero except $\hat{\mu}_{11} = 1$. For

each of the modes, computing the drift and squared diffusion coefficients via (2.14) and (2.15) gives, in case (A1),

$$b_1 = 0, \quad \sigma_1^2 = \frac{3}{7}, \quad b_2 = 0, \quad \sigma_2^2 = \frac{15}{49}. \quad (2.24)$$

In case (A2),

$$b_1 = -\frac{1}{7}, \quad \sigma_1^2 = \frac{3}{7}, \quad b_2 = -\frac{1}{21}, \quad \sigma_2^2 = \frac{15}{49}. \quad (2.25)$$

Note that case (A1) is a single-mode case, whereas case (A2) is a dual-mode case. Roughly speaking, we may infer that, in the former case, in order to perform near optimality, it is necessary to keep the proportions of the work allocated to the different activities close to the fractions given by ξ^A . That is, the policies should be designed to achieve $\Xi^n \Rightarrow \xi^A$. As for case (A2), recall that Z approximates $\hat{W}^n$. Hence, in this case, it is necessary for the work allocation to vary over time in such a way that the aforementioned proportions are close to ξ^L when $\hat{W}^n(t) < z^*$ and to ξ^H when $\hat{W}^n(t) > z^*$. This is, again, only a rough statement; a precise formulation of this behavior appears next.

2.3.3. The Controlled Diffusion. Controlling (2.16) according to the above description results in two different diffusion processes. In the single-mode case, the optimally controlled process Z is given by

$$Z_t^{(1)} = z + b^A t + \sigma^A B_t + L_t^{(1)}, \quad (2.26)$$

with $b^A = b(\xi^A)$ and $\sigma^A = \sigma(\xi^A)$, which is nothing but a reflecting BM with drift b^A and diffusivity σ^A . In the dual-mode case, consider the stochastic differential equation (SDE)

$$Z_t^{(2)} = z + \int_0^t b^*(Z_s^{(2)}) ds + \int_0^t \sigma^*(Z_s^{(2)}) dB_s + L_t^{(2)}, \quad (2.27)$$

where, throughout, we denote

$$b^* = b \circ \varphi^*, \quad \sigma^* = \sigma \circ \varphi^*, \quad \varphi^*(z) = \xi^L \mathbf{1}_{[0, z^*]}(z) + \xi^H \mathbf{1}_{(z^*, \infty)}(z), \quad z \in \mathbb{R}_+. \quad (2.28)$$

For this equation, weak existence and uniqueness of solutions hold, as we shall argue in Lemma 4.1. As a result, there exists a control system for the WCP that behaves exactly as described above, with $\Xi_t = \xi^L$ (respectively, ξ^H) when $Z_t^{(2)} \leq z^*$ ($> z^*$), and, moreover, this description uniquely determines the law of the process $Z^{(2)}$.

As for the asymptotics of the QCP, the preceding discussion and the fact that the system starts empty suggest that in order to achieve the lower bound, the convergence

$$(\hat{X}_p^n, \hat{W}^n) \Rightarrow (0, Z^{(1)}), \quad \text{with } z = 0, \quad (2.29)$$

should hold in the single-mode case, and

$$(\hat{X}_p^n, \hat{W}^n) \Rightarrow (0, Z^{(2)}), \quad \text{with } z = 0, \quad (2.30)$$

in the dual-mode case, where $Z^{(1)}$ is given by (2.26), $(Z^{(2)}, \Xi, L, B)$ is a weak solution to (2.27), and $z = 0$.

2.4. Asymptotic Optimality Results

This section is devoted to the description of several policies that are shown to be AO under different conditions. We have already assumed that the interarrival times of the primitive processes possess finite second moments. Our main results require a stronger assumption.

Assumption 2.8. *There exists $\mathbf{m} > 2$ such that*

$$\max_{i,k} \mathbb{E}[\tilde{a}_i(1)^{\mathbf{m}}] \vee \mathbb{E}[\tilde{u}_{ik}(1)^{\mathbf{m}}] < \infty.$$

Whereas the assumption $\mathbf{m} > 2$ is required for all our results, some of them will require yet a stronger moment assumption—namely, $\mathbf{m} > \mathbf{m}_0$ —where, throughout, we denote

$$\mathbf{m}_0 = \frac{1}{2}(5 + \sqrt{17}).$$

Different policies are proposed in different cases. The distinction between the various cases is based on whether the Single-Mode Condition (2.20) or the Dual-Mode Condition (2.21) holds and, further, for each of the relevant modes (ξ^A in the former case and both ξ^L and ξ^H in the latter), whether the HPC (class p) is the single- or dual-

activity class. Our policies under the single-mode and dual-mode conditions are presented in Definitions 2.11 and 2.12, respectively. These use the $\mathbf{P}$, $\mathbf{T}_1$ and $\mathbf{T}_2$ rules defined in Definition 2.10. The $\mathbf{P}$ rule is the simple priority rule that always give priority to the HPC, whereas the $\mathbf{T}_2$ rule involves a threshold, similar to the policy in Bell and Williams (2001), which prioritizes the HPC, except when the number of HPC customers is below a certain threshold. The $\mathbf{T}_1$ rule also involves a threshold (again, prioritizing the HPC, except when the number of HPC customers is below a certain threshold) and is only used in some cases of dual-mode policies for reasons described in the paragraph following Definition 2.11. Policies under the single-mode condition involve one rule, whereas policies under the dual-mode condition involve one or two rules.

All policies we describe are nonpreemptive; that is, the processing of a job is not interrupted once started. A job is said to be *in the queue* if it is waiting to be served, whereas it is *in the system* if it is either in the queue or being processed. (In what is a bit of an abuse of terminology, we use the term *queue length* to refer to the number in the system.) A server is said to be *available* at a time t if either it has just completed a job or has already been idle at that time.

Some of the policies to be described are defined in terms of a sequence of thresholds, Θ^n , put on the queue length at one of the two buffers. Under Assumption 2.8, $m > 2$. Fix $\bar{a}$ satisfying

$$\frac{1}{2} - \bar{\zeta}(m) < \bar{a} < \frac{1}{2} \quad \text{where} \quad \bar{\zeta}(m) = \begin{cases} \frac{m-2}{4m}, & m \in (2, m_0], \\ \frac{m-2}{4m} \wedge \frac{m^2-5m+2}{2m(3m-2)} = \frac{m^2-5m+2}{2m(3m-2)}, & m \in (m_0, \infty). \end{cases} \quad (2.31)$$

Set the sequence of threshold levels Θ^n and their normalized version $\hat{\Theta}^n$ to

$$\Theta^n = \lceil n\bar{a} \rceil, \quad \hat{\Theta}^n = n^{-1/2}\Theta^n. \quad (2.32)$$

Definition 2.9 (Server Dedicated to/Prioritizes a Class).

1. A server is said to be dedicated to class i at a given time if it acts as follows: if available at that time, it admits a job from class i , provided there is one in the queue, or there is a new class- i arrival, but does not admit a job from the other class.
2. A server is said to prioritize class i at a given time if it acts as follows: if available at that time, it admits a job from class i , provided there is one in the queue; otherwise, it admits a job from the other class, provided there is one in the queue. If the server is idle at that time and there is a new arrival of any class, it admits this arrival unless the other server is dedicated to that class and is free at that time (in which case the other server admits it).

Definition 2.10 ($\mathbf{P}$, $\mathbf{T}_1$, and $\mathbf{T}_2$ rules). Let a mode ξ be given. At any moment in time the single-activity server is dedicated to the dual-activity class.

1. The servers are said to obey the priority rule, abbreviated the $\mathbf{P}$ rule, at a given time if the dual-activity server prioritizes the single-activity class at that time.
2. The servers are said to obey the single-activity class threshold rule, abbreviated the $\mathbf{T}_1$ rule, at a given time if in the n -th system, the dual-activity server prioritizes the single-activity class when the queue length of the single-activity class equals or exceeds Θ^n at that time, and otherwise prioritizes the dual-activity class.
3. The servers are said to obey the dual-activity class threshold rule, abbreviated the $\mathbf{T}_2$ rule, at a given time if in the n -th system, the dual-activity server prioritizes the dual-activity class when the queue length of the dual-activity class equals or exceeds Θ^n at that time, and otherwise prioritizes the single-activity class.

Recall the notations ξ^A , ξ^L , ξ^H , and p from Section 2.3. In the case of a single mode, the following two policies are proposed. (In Definitions 2.11 and 2.12, the text in square brackets is not a part of the definition, but serves to indicate when each policy is to be applied.)

Definition 2.11 (Single-Mode Policies $\mathbf{P}$ and $\mathbf{T}_2$). Let the Single-Mode Condition (2.20) hold.

1. The $\mathbf{P}$ policy [to be applied when $i_1(\xi^A) = p$] is as follows: the servers obey the $\mathbf{P}$ rule corresponding to ξ^A at all times.
2. The $\mathbf{T}_2$ policy [to be applied when $i_2(\xi^A) = p$] is as follows: the servers obey the $\mathbf{T}_2$ rule corresponding to ξ^A at all times.

In the dual-mode case, servers switch between two rules, depending, roughly speaking, on whether the rescaled workload is below or above the switching point z^* . This makes the structure of the policies more complicated. In particular, there is potential loss of capacity in every switching, especially because our policies are non-preemptive. Moreover, the number of switchings grows without bound, as the rescaled workload converges to a diffusion. Therefore, one has to be careful about how to assure that the rule obeyed is updated soon enough after the workload level has crossed the switching point so as not to compromise optimality. The considerations differ

in the different cases and give rise to the use of the T_1 rule, as well as the sampling rules in Definition 2.12. Although it is possible that AO is not too sensitive to these fine details, proving that might be quite hard.

The precise definition requires the use of a state variable called *current mode*, which determines which rule is applicable. This variable is not updated continuously in time about whether $W^n > n^{1/2}z^*$, but only at certain sampling times, as detailed below. The four policies used in the dual-mode case are as follows.

Definition 2.12 (Dual-Mode Policies PP , T_2T_2 , T_1T_2 , and T_2T_1). Let the Dual-Mode Condition (2.21) hold.

1. The PP policy [applied when $i_1(\xi^L) = i_1(\xi^H) = p$] is as follows.
 - a. The workload is sampled at each service completion of the single-activity server. If the workload is below $n^{1/2}z^*$, the current mode is set to $\xi = \xi^L$; otherwise, it is set to $\xi = \xi^H$.
 - b. The servers always obey the P rule with respect to (w.r.t.) the current mode ξ .
2. The T_2T_2 policy [applied when $i_2(\xi^L) = i_2(\xi^H) = p$] is as follows.
 - a. The workload is sampled at each service completion of the single-activity server. If the workload is below $n^{1/2}z^*$, the current mode is set to $\xi = \xi^L$; otherwise, it is set to $\xi = \xi^H$.
 - b. The servers always obey the T_2 rule w.r.t. the current mode ξ .
3. The T_1T_2 policy [applied when $i_1(\xi^L) = i_2(\xi^H) = p$] is as follows.
 - a. The workload is sampled at each arrival and service completion. If the workload is below $n^{1/2}z^*$, the current mode is set to $\xi = \xi^L$; otherwise, it is set to $\xi = \xi^H$.
 - b. Whenever $\xi = \xi^L$, the servers obey the T_1 rule w.r.t. ξ ; whenever $\xi = \xi^H$, they obey the T_2 rule w.r.t. ξ .
4. The T_2T_1 policy [applied when $i_2(\xi^L) = i_1(\xi^H) = p$] is as T_1T_2 , except that the roles of T_1 and T_2 are interchanged.

Our main result states conditions under which each of the six policies just introduced are AO.

Theorem 2.13. Let Assumptions 2.2 and 2.8 hold. In parts 1(b), 2(b), and 2(c), assume moreover that $\mathbf{m} > \mathbf{m}_0$.

1. Assume that the Single-Mode Condition (2.20) holds.
 - a. If $i_1(\xi^A) = p$, then under the P policy, (2.29) holds, and this policy is AO.
 - b. If $i_2(\xi^A) = p$, then under the T_2 policy, (2.29) holds, and this policy is AO.
2. Assume that the Dual-Mode Condition (2.21) holds.
 - a. If $i_1(\xi^L) = i_1(\xi^H) = p$, then under the PP policy, (2.30) holds, and this policy is AO.
 - b. If $i_2(\xi^L) = i_2(\xi^H) = p$, then under the T_2T_2 policy, (2.30) holds, and this policy is AO.
 - c. If $i_1(\xi^L) = i_2(\xi^H) = p$, then under the T_1T_2 policy, (2.30) holds, and this policy is AO.
 - d. If $i_2(\xi^L) = i_1(\xi^H) = p$, then under the T_2T_1 policy, (2.30) holds, and this policy is AO.

The proof of this result is in Section 5. Table 1 summarizes how to determine which of the above six case applies, based on the problem data. We discuss some of the fine details of our construction in the dual-mode case in Section 5.1.3.

Remark 2.14. In Section 5.1, Theorem 5.1 provides more detailed information than (2.29) and (2.30) on the weak limit of the processes involved.

Although the main goal of this paper is to address the case where the LP has multiple solutions, the following result about the case where it has a unique solution is a consequence of our treatment.

Corollary 2.15. Let the “standard” HTC hold; that is, the EHTC with a unique LP solution $(\xi^*, 1)$. Let also Assumption 2.8 hold. In part (b), assume moreover that $\mathbf{m} > \mathbf{m}_0$.

- a. If $i_1(\xi^*) = p$, then under the P policy, (2.29) holds, and this policy is AO.
- b. If $i_2(\xi^*) = p$, then under the T_2 policy, (2.29) holds, and this policy is AO.

Remark 2.16. This result is closely related to the main result of Bell and Williams (2001), which proved AO of a threshold policy under the standard HTC. In a remark in Bell and Williams (2001, p. 622), it was conjectured that the threshold policy without preemption has the same behavior in the heavy traffic limit as the one with preemption. Corollary 2.15 confirms that AO can indeed be achieved by a nonpreemptive threshold policy, although, strictly speaking, it does not prove the precise conjecture from Bell and Williams (2001), as our threshold sizes differ from those of Bell and Williams (2001).

Example 2.17. We continue cases (A1) and (A2) from Example 2.7, specifying now which part of the main result applies in each case. Recall from Example 2.3 that $\alpha_1 = 7$ and $\alpha_2 = 14$. Recall from Example 2.7 that Figure 2 depicts the two modes and the corresponding graphs and that the drift-diffusion pairs are given by (2.24) and (2.25). Case (A1) is a single mode because $b_1 = b_2$, and the active mode is $\xi^A = \xi^{*,2}$ because $\sigma_2 < \sigma_1$. As can be seen in Figure 2 (right), this mode has $i_1(\xi^A) = 2$. Now, assume that the costs are $(h_1, h_2) = (1, 1)$. Then, the class

Table 1. Identifying an AO Policy Assuming LP Solution Multiplicity

Input: the problem data, $(\lambda_i), (\mu_{ik}), (\hat{\lambda}_i), (\hat{\mu}_{ik}), (\sigma_i), (\sigma_{ik}), \gamma, (h_i)$. Output: an AO policy.

1. Find the two modes (i.e., extreme solutions) of the LP (2.7), $\xi^{*,m}$, $m = 1, 2$. In particular, because the LP is assumed to have multiple solutions, μ_{ik} are given as $\alpha_i \beta_k$, and the formulas in Lemma 3.1 express $\xi^{*,m}$, $m = 1, 2$ in terms of α_i and β_k .
2. Compute the drift-diffusion pairs (b_m, σ_m) , $m = 1, 2$. For this, use Equation (2.15).
3. Decide *single-* or *dual-mode* case, and find the active-/lower-/higher-workload modes, as follows. For $(m, m') = (1, 2)$ or $(2, 1)$, if $b_m \leq b_{m'}$ and $\sigma_m \leq \sigma_{m'}$, then this is the single-mode case, with m the active mode: $\xi^A = \xi^{*,m}$. Otherwise, this is the dual-mode case, and if $\sigma_m < \sigma_{m'}$, then m (resp., m') is the lower (higher)-workload mode: $\xi^L = \xi^{*,m}$, $\xi^H = \xi^{*,m'}$.
4. In the dual-mode case, compute the switching point z^* . This is done by numerically solving the HJB Equation (2.18)–(2.19), or, alternatively, Equation (B.3).
5. Determine high-priority class: $p = \arg \max_i h_i \alpha_i$.
6. For $\xi = \xi^A$ (single mode) or $\xi = \xi^L$ and ξ^H (dual mode), denote by $i_1(\xi)$ (resp., $i_2(\xi)$) the class that is assigned one (resp., two) server(s) under mode ξ .
7. Given the number of modes, the modes ξ^A or ξ^L and ξ^H , the indices p and $i_1(\xi)$, $i_2(\xi)$ for the relevant modes ξ , determine which of the six cases of Theorem 2.13 applies.

that maximizes $h_i \alpha_i$ is $p = 2$. Thus, $i_1(\xi^A) = p$, and Theorem 2.13(1(a)) holds. If, instead, $(h_1, h_2) = (3, 1)$, then $p = 1$, and, therefore, Theorem 2.13(1(b)) holds.

In case (A2), we have $b_1 < b_2$ but $\sigma_2 < \sigma_1$; hence, this is a dual-mode case with $\xi^L = \xi^{*,2}$ and $\xi^H = \xi^{*,1}$. By Figure 2, $i_1(\xi^L) = 2$ and $i_1(\xi^H) = 2$. Again, if $(h_1, h_2) = (1, 1)$, then $p = 2$, and we see that Theorem 2.13(2(a)) applies, but if $(h_1, h_2) = (3, 1)$, $p = 1$, and Theorem 2.13(2(b)) applies.

2.5. Some Numerical Results

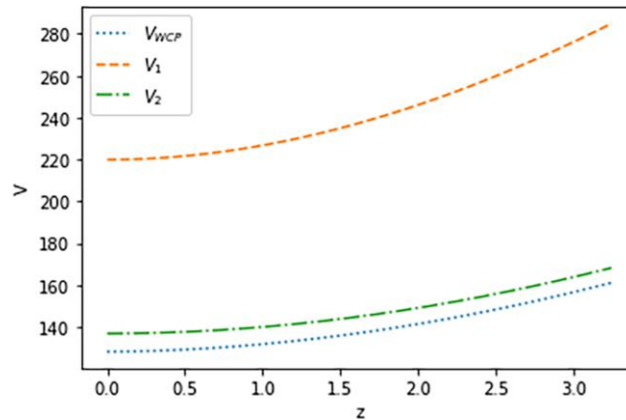
The dual-mode policies that we introduce are admittedly complicated, and the proof of their asymptotic optimality is quite involved. A natural question thus arises here: Is all of this complexity worthwhile? What is the gain from it?

From a purely theoretical/mathematical context, the answer is simple: If the dual-mode policy has a cost that is strictly below that of both single-mode controls, then a nonswitching policy cannot be asymptotically optimal. From a practical viewpoint, this answer falls short. Implementing the dual-mode policy is clearly more complicated than implementing a single-mode policy. A key question thus arises: Is the gain from using a dual-mode policy sufficient to overcome the effort involved in implementing it? The answer to this question is context-dependent and clearly depends on the cost of the effort required to implement the dual-mode control, which we do not include as part of our model. Thus, we do not quite answer this question.

We do, however, partially answer the question of how much larger is the cost of single-mode policies over the optimal switching policy. We do this in two numerical examples. All of the numerical results that we present here are obtained by finding the unique root of Equation (B.3) and then using (B.1) and (B.2). Note that to use (B.3), (B.1), and (B.2), the identity of the two modes needs to be flipped around because in case (A2), we have $b_1 < b_2$, and the analysis in Appendix B requires $b_1 \geq b_2$. The mode identities in our presentation remains as in (A2).

The first numerical results are for the case (A2) introduced in Example 2.7, with $\gamma = 0.01$. In particular, in Figure 4, we plot three curves: $V_1(z)$, $V_2(z)$, and $V_{WCP}(z)$. Here, V_i is the cost function for using mode i , $i = 1, 2$. (The orange/top curve corresponds to V_1 , and the green/middle curve corresponds to V_2 .) It is clear from Figure 4

Figure 4. Value Functions in Case (A2)



that there is a gap between $V_{WCP}(z)$ and both $V_1(z)$ and $V_2(z)$. We define the “gain” from using the optimal switching policy as

$$G := \frac{V_1(0) \wedge V_2(0)}{V_{WCP}(0)}.$$

Thus, $G \geq 1$. For case (A2), $G = 1.07$. If the cost related to implementing the more complicated switching policy is great enough, it may indeed be decided that $G = 1.07$ is not sufficient to overcome this cost.

This raises a more general question: Can we place an a priori upper bound on G ? It is beyond the scope of this paper to answer this question in any precise mathematical manner, but we provide a set of numerical results *suggesting* that the answer may be no: Given any $M < \infty$, it is possible to find a set of parameters $\{(b_i, \sigma_i^2), i = 1, 2\}$ such that $G > M$. We examine a set of parameters loosely related to case (A2). In particular, we fix $b_2 = -1/21$ and $\sigma_2^2 = 15/49$, which are the parameters in case (A2), throughout. We also use $\gamma = 0.01$. We then take six different values for b_1 and set σ_1^2 so that

$$\frac{\sigma_1^2}{|b_1|} = \frac{\sigma_2^2}{|b_2|}.$$

In particular, we take the b_1 values to be $\{-3/21, -6/21, -12/21, -24/21, -48/21, -96/21\}$. Table 2 contains the results of this numerical experiment. Note that $G = 1.26$ when $b_1 = -3/21$ and grows to $G = 5.01$ when $b_1 = -96/21$.

In a very real sense, the situation presented in Figure 4 is a “lucky” one for a system controller who does not want to go through the trouble of implementing the switching policy because the loss is only 7%. The results of Table 2 stand in contrast to that. To see this in graphical form, in Figure 5, we present the same plot as in Figure 4, but corresponding to the parameters in the third row of Table 2. (In this case, V_1 and V_2 have switched places: The green/top curve corresponds to V_2 , and the orange/middle curve corresponds to V_1 .)

3. The LP Under the EHTC

In this section, we prove Lemma 2.1. In addition, in a sequence of three lemmas, we provide an explicit solution to the LP and a criterion for determining whether the two modes are class-switched or server-switched. Finally, Theorem 2.4 is proved. The section is structured as follows. In Section 3.1, we first prove Lemma 2.1(1–3), based mostly on results from Atar et al. (2024). Then, we state Lemmas 3.1 and 3.2, which provide the LP solution, and Lemma 3.3, which is concerned with how the modes are paired. Lemma 2.1(4) is then proved based on Lemma 3.3. In Section 3.2, we prove Lemmas 3.1–3.3 and Theorem 2.4.

3.1. LP-Related Lemmas

Proof of Parts 1–3 of Lemma 2.1.

1. Let $\xi \in \mathcal{S}_{LP}$. It is impossible that $\sum_i \xi_{ik} < 1$ for both $k = 1$ and 2, as this contradicts the EHTC $\rho^* = 1$. Assume, then, that, say, $\sum_i \xi_{i1} < 1$. Then, $\xi_{11} < 1$. Define

$$\tilde{\xi} = \xi + \begin{pmatrix} \varepsilon & -c\varepsilon \\ 0 & 0 \end{pmatrix},$$

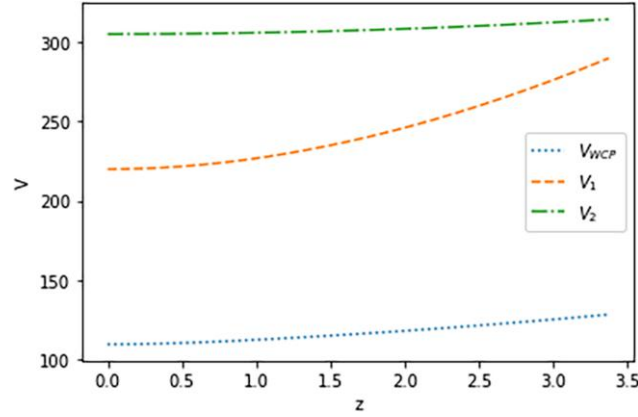
where $c = \mu_{12}^{-1} \mu_{11}$. Then, for $\varepsilon > 0$ small, $\tilde{\xi}$ satisfies (2.7) with $\rho < 1$, which again contradicts the EHTC. This proves Part 1.

2. This follows from Atar et al. (2024, lemma 2.4(4)). Note that uniqueness of the dual problem, which is a standing assumption in Atar et al. (2024), is not used in the proof of this statement.

Table 2. Gain from Using Switching Policy

b_1	σ_1^2	z^*	G	V_{WCP}
$-3/21$	$45/49$	1.91	1.26	174
$-6/21$	$90/49$	1.49	1.56	141
$-12/21$	$180/49$	1.13	2.01	109
$-24/21$	$360/49$	0.84	2.67	82
$-48/21$	$720/49$	0.61	3.63	61
$-96/21$	$1,440/49$	0.44	5.01	44

Figure 5. Value Functions for Row 3 of Table 2



3. This statement follows from Atar et al. (2024, lemma 2.3(1)) and Atar et al. (2024, lemma 2.4(4)), where, again, the uniqueness of the dual is not used. $\square$

The following lemma computes the two modes.

Lemma 3.1. *Let the EHTCM hold. Then,*

$$\sum_i \frac{\lambda_i}{\alpha_i} = \sum_k \beta_k = 1, \quad (3.1)$$

where the last equality merely expresses the normalization convention mentioned earlier in Section 2.2. Moreover, any $\xi \in S_{LP}$ is determined by its entry ξ_{11} via

$$\xi = \begin{pmatrix} \xi_{11} & \frac{\lambda_1}{\alpha_1\beta_2} - \frac{\beta_1}{\beta_2}\xi_{11} \\ 1 - \xi_{11} & 1 - \frac{\lambda_1}{\alpha_1\beta_2} + \frac{\beta_1}{\beta_2}\xi_{11} \end{pmatrix}. \quad (3.2)$$

The two modes $\xi^{*,1}$ and $\xi^{*,2}$ can be expressed by (3.2) with ξ_{11} given by

$$\xi_{11}^{*,1} = \max\left(0, \frac{\lambda_1}{\alpha_1\beta_1} - \frac{\beta_2}{\beta_1}\right), \quad \xi_{11}^{*,2} = \min\left(\frac{\lambda_1}{\alpha_1\beta_1}, 1\right). \quad (3.3)$$

Recall that under the nondegeneracy condition, for any mode, there is a unique relabeling of classes and servers, which transforms it to a canonical form. The following result shows that both modes, once put in canonical form, are given by the same formula.

Lemma 3.2. *Let Assumption 2.2 (EHTCM and nondegeneracy) hold. Fix $m \in \{1, 2\}$. Relabel classes and servers so that $\xi^{*,m}$ is in canonical form. Then, $\lambda_1 > \alpha_1\beta_1$ and*

$$\xi^{*,m} = \begin{pmatrix} 1 & \frac{\lambda_1}{\alpha_1\beta_2} - \frac{\beta_1}{\beta_2} \\ 0 & \frac{\lambda_2}{\alpha_2\beta_2} \end{pmatrix}. \quad (3.4)$$

In particular, if $\xi, \xi' \in S_{LP}$ and there are i, k such that $\xi_{ik} = \xi'_{ik} = 0$, then $\xi = \xi'$.

Lemma 2.1(4), which is yet to be proved, states that under the nondegeneracy condition, the two modes must be either class- or server-switched. The following lemma contains this result and, in addition, provides a criterion for distinguishing between these cases. We will say that the *class-switching condition* holds if

$$\max_i \frac{\lambda_i}{\alpha_i} < \max_k \beta_k, \quad (3.5)$$

and the *server-switching condition* holds if

$$\max_i \frac{\lambda_i}{\alpha_i} > \max_k \beta_k. \quad (3.6)$$

Lemma 3.3. *Let Assumption 2.2 hold. Then, both modes are nondegenerate. Moreover, under the Class-Switching Condition (3.5), the modes are class-switched (as, for example, in Figure 1, (a) and (b)), and under the Server-Switching Condition (3.6), they are server-switched (as, for example, in Figure 1, (a) and (c)).*

Proof of Part 4 of Lemma 2.1. The statement is contained in Lemma 3.3. $\square$

Remark 3.4. Note that cases 2(c) and 2(d) of Theorem 2.13 correspond to class-switched modes (for example, under case 2(c), one has $i_1(\xi^L) = i_2(\xi^H) = p$; hence, the nonbasic activity must have switched from class p to class q when moving from ξ^L to ξ^H). By Lemma 3.3, this occurs under (3.5). Also note that in both cases, the proposed policies apply a different rule for the lower and upper workload modes. On the other hand, cases 2(a) and 2(b) of Theorem 2.13 correspond to server-switching and hold under (3.6), and our policies are such that the same rule is used for the lower and upper workload modes.

3.2. Proof of Lemmas 3.1–3.3 and Theorem 2.4

Proof of Lemma 3.1. In view of Lemma 2.1(1), every solution ξ is column-stochastic, and, recalling $\mu_{ik} = \alpha_i \beta_k$, must satisfy

$$\begin{aligned}\xi_{11}\beta_1 + \xi_{12}\beta_2 &= \frac{\lambda_1}{\alpha_1}, \\ \xi_{21}\beta_1 + \xi_{22}\beta_2 &= \frac{\lambda_2}{\alpha_2}, \\ \xi_{11} + \xi_{21} &= 1, \\ \xi_{12} + \xi_{22} &= 1, \\ \xi_{i,k} &\geq 0, \quad i, k \in \{1, 2\}.\end{aligned}\tag{3.7}$$

Identity (3.1) follows.

Next, because the Expression (3.2) is also column-stochastic, proving that any solution ξ is determined by ξ_{11} as in (3.2) amounts to proving that ξ_{12} is given as in (3.2). This follows from the first line in (3.7).

It remains to prove (3.3). By the expression just obtained for ξ_{12} , it follows that as long as

$$\xi_{11} \geq \frac{\lambda_1}{\alpha_1\beta_1} - \frac{\beta_2}{\beta_1},$$

we obtain $\xi_{12} \leq 1$. Clearly, in addition, $\xi_{11} \geq 0$ must hold. Similarly, as long as

$$\xi_{11} \leq \frac{\lambda_1}{\alpha_1\beta_1},$$

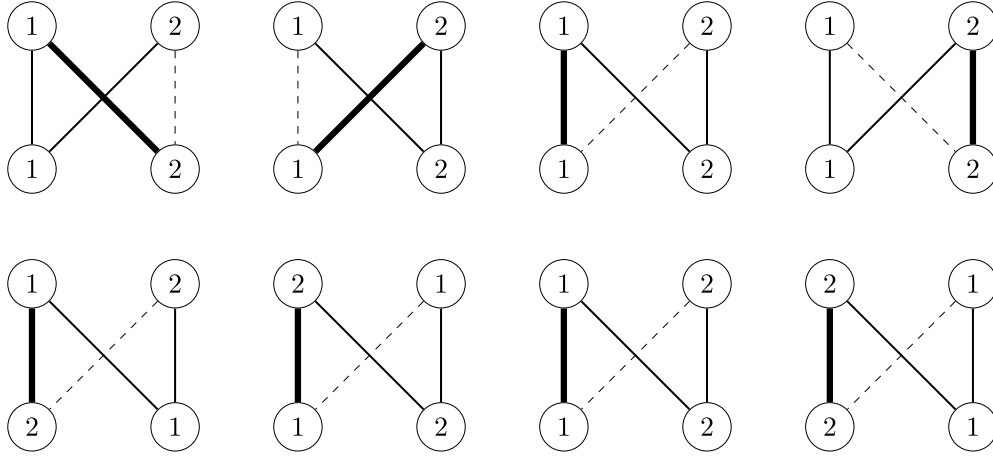
we obtain $\xi_{12} \geq 0$, and, in addition, $\xi_{11} \leq 1$ must hold. As a result, it is necessary that

$$\xi_{11} \in \left[\max\left(0, \frac{\lambda_1}{\alpha_1\beta_1} - \frac{\beta_2}{\beta_1}\right), \min\left(\frac{\lambda_1}{\alpha_1\beta_1}, 1\right) \right].\tag{3.8}$$

Moreover, setting ξ_{11} to each of the two endpoints of the interval indicated in (3.8) and letting ξ be the corresponding expression from (3.2) gives rise to a solution satisfying all of (3.7), as can be checked directly. Because by (3.2) a solution ξ is an affine function of its entry ξ_{11} , these two endpoints correspond to the two extreme points of $\mathcal{S}_{LP}$; that is, to the two modes $\xi^{*,1}$, $\xi^{*,2}$. This proves the lemma. $\square$

Proof of Lemma 3.2. Note that Relations (3.7) are invariant to relabeling of classes and servers. Hence, so is Relation (3.2), which was derived solely from (3.7). Let m be given and assume a relabeling has been performed to put $\xi^{*,m}$ in canonical form. Then, $\xi^{*,m}$ satisfies (3.2) with its first column given by $(1, 0)^T$. Consequently $\xi_{11}^{*,m} = 1$. Substituting $\xi_{11}^{*,m} = 1$ into (3.2) proves (3.4). Because under the nondegeneracy assumption, there can be only one zero entry, in (3.4), we have $\xi_{12}^{*,m} > 0$. Hence, $\lambda_1 > \alpha_1\beta_1$. The final assertion follows from (3.4) using, again, the fact that there can be at most one zero entry $\square$

Figure 6. Top: Four Possible Graphs of Basic Activities; Bottom: Corresponding Relabelings for Canonical Form



The four possible graphs and their relabelings are described in Figure 6. Namely, if (i', k) is the nonbasic activity in $\xi^{*,m}$, then defining

$$\tilde{\xi}_{11}^{*,m} = \xi_{ik}^{*,m} = 1,$$

$$\tilde{\xi}_{22}^{*,m} = \xi_{i'k'}^{*,m},$$

$$\tilde{\xi}_{21}^{*,m} = \xi_{i'k}^{*,m} = 0, \text{ and}$$

$$\tilde{\xi}_{12}^{*,m} = \xi_{ik'}^{*,m},$$

$\tilde{\xi}^{*,m}$ is obtained from $\xi^{*,m}$ upon relabeling in the form of an “N.”

Proof of Lemma 3.3. The nondegeneracy of both modes follows from Lemma 3.2.

Next, let the Class-Switching Condition (3.5) hold. Because of (3.1),

$$\max_k \beta_k > \max_i \frac{\lambda_i}{\alpha_i} \geq \min_i \frac{\lambda_i}{\alpha_i} > \min_k \beta_k. \quad (3.9)$$

Recall from the proof of Lemma 3.1 that the two endpoints of the interval defined in (3.8) correspond to the two modes. Consider the right endpoint. If the minimum in expression in (3.8) is one, then, by (3.2), the nonbasic activity in that mode is $(2, 1)$, and, moreover, $\frac{\lambda_1}{\alpha_1} > \beta_1$. By (3.1), this gives $\frac{\lambda_2}{\alpha_2} < \beta_2$. In view of (3.9), this gives

$$\beta_2 > \max_i \frac{\lambda_i}{\alpha_i}.$$

Hence, the maximum in (3.8) is zero. By (3.2), this shows that the nonbasic activity in the other mode is $(1, 1)$. If, on the other hand, the minimum in (3.8) is $\frac{\lambda_1}{\alpha_1 \beta_1}$, then $\frac{\lambda_1}{\alpha_1} < \beta_1$, and the nonbasic activity in the corresponding mode is $(1, 2)$. Similarly, by (3.9),

$$\min_i \frac{\lambda_i}{\alpha_i} > \beta_2.$$

This means the maximum in (3.8) is not zero. The nonbasic activity in the other mode is then $(2, 2)$. In both cases, the two modes form a class-switched pair as claimed.

Consider now the Server-Switching Condition (3.6). Because of (3.1),

$$\max_i \frac{\lambda_i}{\alpha_i} > \max_k \beta_k \geq \min_k \beta_k > \min_i \frac{\lambda_i}{\alpha_i}. \quad (3.10)$$

If the minimum in (3.8) is one, then $\frac{\lambda_1}{\alpha_1} > \beta_1$ and the nonbasic activity in one mode is $(2, 1)$. By (3.10), $\frac{\lambda_1}{\alpha_1} > \beta_2$. Hence, the maximum in (3.8) is not zero, and the nonbasic activity in the other mode is $(2, 2)$. Finally, if the minimum in (3.8) is not one, then $\frac{\lambda_1}{\alpha_1} < \beta_1$, and the nonbasic activity in one mode is $(1, 2)$. By (3.10), $\frac{\lambda_1}{\alpha_1} < \beta_2$. Hence, the

maximum in (3.8) is zero and the nonbasic activity in the other mode is (1, 1). In both cases, the two of modes forms a server-switched pair. $\square$

Proof of Theorem 2.4. This lower bound is precisely the one stated in Atar et al. (2024, theorem 2.6), when specialized to the 2×2 PSS. To prove that it is valid, we must verify that the standing assumption of Atar et al. (2024), namely, Atar et al. (2024, assumption 2.2), holds.

First, Atar et al. (2024, assumption 2.2.1), which states that the EHTC holds, is valid because of our Assumption 2.2.1. Next, Atar et al. (2024, assumption 2.2.2), when translated to the notation of this paper, states that every $\xi \in \mathcal{S}_{LP}$ is column-stochastic. This holds by our Lemma 2.1.1. It remains to show that Atar et al. (2024, assumption 2.2.3), which states that the dual of (2.7) has a unique solution, is satisfied.

To this end, we shall adopt in the remainder of this proof some notation and terminology from Atar et al. (2024). By Lemma 3.2, the nonbasic activity is different in both modes $\xi^{*,1}, \xi^{*,2}$, for else one would have $\xi^{*,1} = \xi^{*,2}$, contradicting the EHTCM. As a result, with the terminology introduced in Atar et al. (2024, section 2), all the activities are potentially basic. Using strict complementary slackness ((36) in Schrijver 1986, chapter 7) in the same way as in Atar et al. (2024, lemma 2.3.4), any (y, z) solution of the dual satisfies

$$y_i = \mu_j z_k, \quad i, k \in \{1, 2\}, j = (i, k).$$

It follows that $y_1 y_2 = \mu_{11} \mu_{21} z_1^2 = \mu_{12} \mu_{22} z_2^2$. By positivity of both z_k , this gives $z_1 = c z_2$ for some constant $c > 0$. Moreover, one of the constraints of the dual problem Atar et al. (2024, equation (2.8)) is $z_1 + z_2 = 1$. Therefore, (z_k) are uniquely determined. As a consequence, so are (y_i) , which shows that the dual problem has at most one solution. The existence of a dual solution follows from the EHTC, as shown in Atar et al. (2024, lemma 2.4.2). This completes the verification of Atar et al. (2024, assumption 2.2). $\square$

4. The WCP and HJB Equation

In this section, Proposition 2.6 is proved. Lemma 4.1, which is used to prove it, contains two additional results: An identification of optimal control systems for the WCP, and weak uniqueness of solutions to the SDE (2.27), both needed for the weak convergence proofs in Section 5.

Proof of Part 1 of Proposition 2.6. This is a special case of Atar et al. (2024, proposition 2.5). We comment that the fact that V_{WCP} is a classical solution to (2.18) and (2.19) has been established already in Sheng (1978). However, uniqueness of solutions is not covered there. $\square$

In what follows, u always denotes the unique solution to (2.18) and (2.19).

In the following lemma, parts 1 and 2 are largely based on results from Sheng (1978). For completeness, we have included details on the construction from Sheng (1978) (as corrected in Sheng 1981) in Appendix B.

Lemma 4.1.

1. (Optimality in the single-mode case). Assume (2.20) and recall that in this case $\xi^A = \xi^{*,m}$ for m as in (2.20). Then, with u as above, Equation (2.22) of Proposition 2.6 holds. Moreover, $J_{WCP}(z, \mathfrak{S}^{(1)}) = V_{WCP}(z)$ for the admissible control system $\mathfrak{S}^{(1)} = (\Omega, \mathcal{F}, \{\mathcal{F}_t\}, \mathbb{P}, B, \Xi, Z^{(1)}, L^{(1)})$, where $(Z^{(1)}, L^{(1)}, B)$ is the RBM from (2.26) (assumed to be constructed on the original probability space), $\mathcal{F}_t = \sigma\{B_s : s \in [0, t]\}$ and $\Xi_t = \xi^A$.
2. (Optimality in the dual-mode case). Assume (2.21) and recall that in this case $(\xi^L, \xi^H) = (\xi^{*,m'}, \xi^{*,m})$. Then, there exists a switching point $z^* \in (0, \infty)$ such that (2.23) of Proposition 2.6 holds. Moreover, SDE (2.27) possesses a weak solution $(\Omega', \mathcal{F}', \{\mathcal{F}'_t\}, \mathbb{P}', B, Z^{(2)}, L^{(2)})$. Furthermore, one has $J_{WCP}(z, \mathfrak{S}^{(2)}) = V_{WCP}(z)$ for the admissible control system defined by $\mathfrak{S}^{(2)} = (\Omega', \mathcal{F}', \{\mathcal{F}'_t\}, \mathbb{P}', B, \Xi, Z^{(2)}, L^{(2)})$, where $\Xi_t = \varphi^*(Z_t^{(2)})$.
3. Weak uniqueness holds for solutions to SDE (2.27).

Proof of Parts 2 and 3 of Proposition 2.6. These results are contained in parts 1 and 2 of Lemma 4.1. $\square$

For Markov control problems, a map from the state space to the control action space is often called a *stationary (feedback) control policy*, or a *policy* for short. For our WCP, a policy is thus a measurable map $\bar{\xi} : \mathbb{R}_+ \rightarrow \mathcal{S}_{LP}$. This term is used in Sheng (1978), and we adopt it in the next proof. Parts 1 and 2 take full advantage of several results from Sheng (1978), where a classical solution to Equations (2.18)–(2.19) is constructed, and a description of an optimal policy is provided.

Proof of Parts 1 and 2 of Lemma 4.1. The Single-Mode Condition (2.20) corresponds to Sheng (1978, section 5.3, equations (41) and (42)). The Dual-Mode Condition (2.21) corresponds to the complementary case. It is stated in Sheng (1978, theorem 1, chapter 4) that a policy is optimal if and only if the cost associated to it is C^2 and satisfies Sheng (1978, equation (14), chapter 4), which is the HJB Equation (2.18)–(2.19) in our notation. However, the

equation studied there is more general, and in order to reduce it to our (2.18)–(2.19), one must take the switching costs to vanish (by setting the expression $K = 0$) and the reflection-absorption parameter to correspond to reflection only (by setting $\lambda = 1$).

Under the single-mode condition, the policy constructed has the simple form

$$\bar{\xi}(z) = \xi^{*,m}, \quad z \in \mathbb{R}_+,$$

with the same m as in (2.20). It is shown that it is an optimal policy by computing the cost associated it and showing that it solves the HJB equation with its boundary conditions. The fact that (2.22) holds in this case is a direct consequence of the fact that the cost associated to the single-mode policy indeed solves the HJB equation. This completes the proof of part 1.

As for part 2, the claim that the SDE (2.27) possesses a weak solution is proved in Atar et al. (2024, lemma 4.1). Under the dual-mode case, the policy constructed in Sheng (1978) is

$$\bar{\xi}_{z^*}(z) = \xi^{*,m'} \mathbb{1}_{z \leq z^*} + \xi^{*,m} \mathbb{1}_{z > z^*}, \quad z \in \mathbb{R}_+, \quad (4.1)$$

with the same m and m' as in (2.21). By Sheng (1978, theorem 4, chapter 3), any policy of this form gives rise to a cost function that is C^1 in all of $(0, \infty)$ and C^2 in $(0, \infty) \setminus \{z^*\}$. The value of z^* that leads to an optimal policy is found by the principle of smooth fit, namely by requiring that the second derivative is continuous also at z^* . The equation that states this smoothness condition, Sheng (1978, (61), chapter 5) (as corrected in Sheng 1981, (3.21)), turns out to have a unique solution $z^* \in (0, \infty)$. The solution to the HJB equation is then the cost associated to $\bar{\xi}_{z^*}$ with that value of z^* . The statement of part 2, by which the corresponding control system is optimal for the WCP, therefore follows from Sheng (1978, theorem 1, chapter 4). Once again, the very fact that this function solves the HJB equation implies that Equation (2.23) holds in this case. $\square$

Proof of Part 3 of Lemma 4.1. We slightly simplify the notation by writing (2.27)–(2.28) as

$$Z_t = z + \int_0^t b^*(Z_s) ds + \int_0^t \sigma^*(Z_s) dB_s + L_t,$$

where, with $a^*(x) = \sigma^*(x)^2$, there exist constants $b_0, b_1 \in \mathbb{R}$ and $a_0 > 0, a_1 > 0$, such that

$$b^*(x) = \begin{cases} b_0, & x \leq z^*, \\ b_1, & x > z^*, \end{cases} \quad a^*(x) = \begin{cases} a_0, & x \leq z^*, \\ a_1, & x > z^*. \end{cases}$$

In the remainder of this proof, a^* and b^* are written as a and b . Let $\mathcal{A}$ be the operator

$$\mathcal{A}f(x) := b(x)f'(x) + \frac{1}{2}a(x)f''(x),$$

on the domain $\mathcal{D}(\mathcal{A}) := \{f \in C_0^\infty[0, \infty) : f'(0) = 0\}$, where $C_0^\infty[0, \infty)$ denotes the set of compactly supported members of $C^\infty[0, \infty)$. If $f \in \mathcal{D}(\mathcal{A})$ and (Z, L, B) is a weak solution to the SDE, then, by Ito's formula, the boundary property of L and the boundary condition $f'(0) = 0$, the process $f(Z_t) - \int_0^t \mathcal{A}f(Z_s) ds$ is a martingale. Therefore, it follows from the existence result in part 2 of the lemma that, for every probability distribution ν on $[0, \infty)$, there exists a solution of the $D_{[0, \infty)}[0, \infty)$ martingale problem for $(\mathcal{A}, \nu)$. We will use the definition of the stopped martingale problem of Ethier and Kurtz (1986, chapter 4, section 6). Then, for every probability distribution ν on $[0, \infty)$, there exists a solution of the stopped martingale problem for $(\mathcal{A}, \nu, [0, \frac{2}{3}z^*))$ and of the stopped martingale problem for $(\mathcal{A}, \nu, (\frac{1}{3}z^*, \infty))$.

Define the operators $\mathcal{A}_0$ and $\mathcal{A}_1$ in the following way:

$$\begin{aligned} \mathcal{A}_0 f(x) &:= b_0 f'(x) + \frac{1}{2} a_0 f''(x), & \text{on the domain } \mathcal{D}(\mathcal{A}_0) &:= \mathcal{D}(\mathcal{A}) \\ \mathcal{A}_1 f(x) &:= b(x) f'(x) + \frac{1}{2} a(x) f''(x), & \text{on the domain } \mathcal{D}(\mathcal{A}_1) &:= C_0^\infty(\mathbb{R}). \end{aligned}$$

The $D_{[0, \infty)}[0, \infty)$ martingale problem for $\mathcal{A}_0$ is well posed (for instance by corollary 8.1.2, theorem 4.4.1, and proposition 4.3.1 of Ethier and Kurtz 1986). Therefore, for every probability distribution ν on $[0, \infty)$, there exists a unique solution of the stopped martingale problem for $(\mathcal{A}_0, \nu, [0, \frac{2}{3}z^*))$ by theorem 4.6.1 of Ethier and Kurtz (1986). Because, for every $f \in \mathcal{D}(\mathcal{A}_0) = \mathcal{D}(\mathcal{A})$,

$$\mathcal{A}_0 f|_{[0, \frac{2}{3}z^*]} = \mathcal{A}f|_{[0, \frac{2}{3}z^*]},$$

every solution of the stopped martingale problem for $(\mathcal{A}, \nu, [0, \frac{2}{3}z^*))$ is also a solution of the stopped martingale

problem for $(\mathcal{A}_0, \nu, [0, \frac{2}{3}z^*))$; therefore, the solution of the stopped martingale problem for $(\mathcal{A}, \nu, [0, \frac{2}{3}z^*))$ is unique.

Next, the $D_{\mathbb{R}}[0, \infty)$ martingale problem for $\mathcal{A}_1$ is well posed by exercise 7.3.3 of Stroock and Varadhan (2006) (it is shown there that the $C_{\mathbb{R}}[0, \infty)$ martingale problem for $\mathcal{A}_1$ is well posed, but every solution of the $D_{\mathbb{R}}[0, \infty)$ martingale problem for $\mathcal{A}_1$ has paths in $C_{\mathbb{R}}[0, \infty)$ almost surely). Therefore, for every probability distribution ν on $[0, \infty)$, there exists a unique solution of the stopped martingale problem for $(\mathcal{A}_1, \nu, (\frac{1}{3}z^*, \infty))$ by theorem 4.6.1 of Ethier and Kurtz (1986). Because for every $f_1 = \mathcal{D}(\mathcal{A}_1)$ there exists $f \in \mathcal{D}(\mathcal{A})$ such that

$$f|_{[\frac{1}{3}z^*, \infty)} = f_1|_{[\frac{1}{3}z^*, \infty)}, \quad \mathcal{A}f|_{[\frac{1}{3}z^*, \infty)} = \mathcal{A}_1 f_1|_{[\frac{1}{3}z^*, \infty)},$$

every solution of the stopped martingale problem for $(\mathcal{A}, \nu, (\frac{1}{3}z^*, \infty))$ is also a solution of the stopped martingale problem for $(\mathcal{A}_1, \nu, (\frac{1}{3}z^*, \infty))$; therefore, the solution of the stopped martingale problem for $(\mathcal{A}, \nu, (\frac{1}{3}z^*, \infty))$ is unique.

Now, one can apply theorem 4.6.2 in Ethier and Kurtz (1986) with $U_1 := [0, \frac{2}{3}z^*)$, $U_k := (\frac{1}{3}z^*, \infty)$ for $k \geq 2$, to conclude that, for every probability distribution ν on $[0, \infty)$, the solution of the $D_{[0, \infty)}[0, \infty)$ martingale problem for $(\mathcal{A}, \nu)$ is unique.

Finally, because, as already mentioned, every solution to the SDE is a solution to the martingale problem, the weak uniqueness of solutions to the SDE follows. $\square$

Remark 4.2 (Symmetry Conditions). As Lemma 4.1 states, (2.20) is a necessary and sufficient condition for the optimal control of the WCP to not switch between modes. It is natural to ask whether it can be determined if the single- or the dual-mode condition holds under various symmetry conditions. Specifically, consider

$$\frac{\hat{\mu}_{1k}}{\alpha_1} = \frac{\hat{\mu}_{2k}}{\alpha_2}, \quad k = 1, 2, \quad (4.2)$$

$$\frac{\hat{\mu}_{i1}}{\beta_1} = \frac{\hat{\mu}_{i2}}{\beta_2}, \quad i = 1, 2, \quad (4.3)$$

$$C_{S_{i1}} = C_{S_{i2}}, \quad i = 1, 2. \quad (4.4)$$

As it turns out, each of the above three conditions is sufficient for (2.20). This is proved in Lemma C.1 in Appendix C. As a result, each of these conditions is sufficient for nonswitching. These conditions can arise naturally and occur in certain cases in the literature.

5. Asymptotic Optimality

In this section, we prove Theorem 2.13. Toward this, an important intermediate goal is to establish a weak convergence result, stated in Theorem 5.1. The proofs of both theorems rely on four main steps stated in Propositions 5.3–5.6.

5.1. Weak Convergence

5.1.1. Statement of Weak Convergence Result. We will adopt the following convention regarding the six cases listed in Theorem 2.13. When a statement is said to hold in a certain case of Theorem 2.13, it is meant that the assumptions, as well as the policy, specified in this case are in force. For example, saying that a certain statement holds in case 1(a) of Theorem 2.13 means that it holds when the Single-Mode Condition (2.20) and the condition $i_1(\xi^A) = p$ hold and the $\mathbf{P}$ policy is applied, and, moreover, the weaker moment assumption $\mathbf{m} > 2$ is assumed (but $\mathbf{m} > \mathbf{m}_0$ in case 1(b), for example). When a claim is stated without specifying a case, it is meant that it holds in each one of the six.

In addition to the rescaled processes already defined, the weak convergence results will be concerned with additional rescaled processes, namely,

$$\hat{A}_i^n(t) = n^{-1/2}(A_i^n(t) - \lambda_i^n t), \quad \hat{S}_{ik}^n(t) = n^{-1/2}(S_{ik}^n(t) - \mu_{ik}^n t), \quad (5.1)$$

$$\hat{I}_k^n(t) = n^{1/2}I_k^n(t), \quad \hat{L}^n(t) = \beta_1 \hat{I}_1^n(t) + \beta_2 \hat{I}_2^n(t). \quad (5.2)$$

With this notation, the Balance Equation (2.3) for X^n translates under scaling to

$$\hat{X}_i^n(t) = \hat{A}_i^n(t) + n^{-1/2} \lambda_i^n t - \sum_k \hat{S}_{ik}^n(T_{ik}^n(t)) - \sum_k n^{-1/2} \mu_{ik}^n T_{ik}^n(t). \quad (5.3)$$

$$= \hat{A}_i^n(t) - \sum_k \hat{S}_{ik}^n(T_{ik}^n(t)) + (n^{1/2} \lambda_i + \hat{\lambda}_i^n) t - \sum_k T_{ik}^n(t) (n^{1/2} \alpha_i \beta_k + \hat{\mu}_{ik}^n). \quad (5.4)$$

Thus, if we denote

$$\hat{F}_t^n = \sum_i \frac{\hat{A}_i^n(t)}{\alpha_i} - \sum_{ik} \frac{\hat{S}_{ik}^n(T_{ik}^n(t))}{\alpha_i} + \sum_i \frac{\hat{\lambda}_i^n t}{\alpha_i} - \sum_{ik} \frac{\hat{\mu}_{ik}^n T_{ik}^n(t)}{\alpha_i}, \quad (5.5)$$

we obtain the following representation for the workload:

$$\hat{W}_t^n = \hat{F}_t^n + \hat{L}^n(t). \quad (5.6)$$

The convergence of the rescaled primitives is a direct consequence of the central limit theorem for renewal processes (Billingsley 1999, section 17). Namely, the tuple $(\hat{A}_i^n, \hat{S}_{ik}^n)$ converges to what we denote by (A_i, S_{ik}) , comprising six mutually independent BMs with zero drift and diffusivity given by the constants $\lambda_i^{1/2} C_{A_i}$ and $\mu_{ik}^{1/2} C_{S_{ik}}$, which in Section 2.3, we have denoted by $\sigma_{A,i}$ and $\sigma_{S,ik}$, respectively.

Theorem 5.1. *Let the assumptions of Theorem 2.13 hold. Then, $(T^n, \hat{W}^n, \hat{L}^n, \hat{X}^n) \Rightarrow (T, W, L, X)$, where the latter is defined as follows.*

1. In cases 1, (a) and (b) of Theorem 2.13, (W, L) is the RBM and boundary term given by (2.26) and initial condition $z = 0$, and $T_t = \xi^A t$ for all t .
2. In cases 2, (a)–(d) of Theorem 2.13, (W, L) is the (unique in law) weak solution to the SDE (2.27) with initial condition $z = 0$, and letting $\Xi_t = \varphi^*(W_t)$, one has $T_t = \int_0^t \Xi_s ds$ for $t \geq 0$.
3. In all cases, $X_p = 0$ and $X_q = \alpha_q W$.

The significance of this result is that it implies that, under each of the proposed policies, limits of the processes exist and form admissible control systems for the WCP, which are, in view of Lemmas 4.1.1–2, optimal.

We next present a lemma from Atar et al. (2024) concerning limits of $(\hat{A}^n, \hat{S}^n, T^n, \hat{F}^n)$ under general sequences of policies. The Skorohod map, $\Gamma : D_{\mathbb{R}}[0, \infty) \rightarrow D_{\mathbb{R}_+}[0, \infty) \times D_{\mathbb{R}^+}^*[0, \infty)$, takes a function ψ to a pair (φ, η) , where

$$\varphi(t) = \psi(t) + \eta(t), \quad \eta(t) = \sup_{0 \leq s \leq t} \psi(s)^-, \quad t \geq 0. \quad (5.7)$$

The corresponding maps $\psi \mapsto \varphi$ and $\psi \mapsto \eta$ are denoted by Γ_1 and Γ_2 , respectively.

Lemma 5.2. *Let $\{T^n\}$, $T^n \in \mathcal{A}^n$ be any sequence of admissible controls for the QCP for which $\limsup_n \hat{J}^n(T^n) < \infty$. Then, the following conclusions hold. The sequence $(\hat{A}^n, \hat{S}^n, T^n)$ is C-tight. Along any convergent subsequence where $(\hat{A}^n, \hat{S}^n, T^n) \Rightarrow (A, S, T)$, one has*

$$(\hat{A}^n, \hat{S}^n, T^n, \hat{F}^n) \Rightarrow (A, S, T, F),$$

where F is defined in (5.8). There exist on $(\Omega, \mathcal{F})$ processes (B, Ξ, Z', L') and a filtration $\{\mathcal{H}_t\}$ such that the tuple $\mathfrak{S} = (\Omega, \mathcal{F}, \{\mathcal{H}_t\}, \mathbb{P}, B, \Xi, Z', L')$ forms an admissible control system for the WCP with initial condition 0. These processes satisfy the relations

$$\begin{aligned} T &= \int_0^\cdot \Xi_s ds, \quad (Z', L') = \Gamma[F], \\ F_t &= \sum_i \frac{A_i(t) + \hat{\lambda}_i t}{\alpha_i} - \sum_{i,k} \frac{S_{ik}(T_{ik}(t)) + \hat{\mu}_{ik} T_{ik}(t)}{\alpha_i} \\ &= \int_0^t b(\Xi_s) ds + \int_0^t \sigma(\Xi_s) dB_s. \end{aligned} \quad (5.8)$$

Proof. This is the content of Atar et al. (2024, lemmas 5.1 and 5.5). $\square$

This lemma is our starting point for proving Theorem 5.1. Whereas it relates limits of processes associated with the QCP to an admissible control system for the WCP, note that it does not make any claim regarding $\hat{X}^n$ or $\hat{W}^n$ and, hence, by itself is not sufficient to relate the prelimit cost (defined in terms of $\hat{X}^n$) to the WCP cost. In

particular, the pair (Z', L') need not be the weak limit of $(\hat{W}^n, \hat{L}^n)$. To proceed, one must show that under the proposed policies, along the sequence specified in Lemma 5.2, one has

$$(\hat{A}^n, \hat{S}^n, T^n, \hat{F}^n, \hat{W}^n, \hat{L}^n) \Rightarrow (A, S, T, F, W, L), \quad (5.9)$$

where $(W, L) = (Z', L')$. Once this is achieved, the lemma guarantees that the weak limit (W, L) satisfies

$$W_t = \int_0^t b(\Xi_s) ds + \int_0^t \sigma(\Xi_s) dB_s + L_t,$$

with $\int_{[0, \infty)} W_t dL_t = 0$, assuring that the limit tuple indeed forms an admissible system for the WCP. To go from here to the statements made in Theorem 5.1, one further needs to show that (W, L) are as given in (2.26) or (2.27) and that $X_p = 0$.

The propositions presented below address these issues as follows. Proposition 5.3 provides uniform integrability required to eventually deduce Theorem 2.13 from Theorem 5.1 and, in addition, ensures that the prelimit cost remains bounded so Lemma 5.2 may be applied. Proposition 5.4 shows that $\hat{X}_p^n \rightarrow 0$ in probability. Proposition 5.5 states precisely what is needed to attain (5.9). Finally, Proposition 5.6 implies that (2.27) holds in the dual-mode case.

5.1.2. Main Steps Toward Weak Convergence. Throughout what follows, the assumptions of Theorem 2.13 are in force. The four main steps required to achieve weak convergence and, later, Theorem 2.13, are as follows.

Proposition 5.3. *There exists $\varepsilon_0 > 0$ such that*

$$\limsup_n \mathbb{E} \int_0^\infty e^{-\gamma t} (\hat{H}_t^n)^{1+\varepsilon_0} dt < \infty.$$

As already mentioned, this uniform integrability result will allow us to deduce convergence of the costs, as stated in Theorem 2.13, from the convergence stated in Theorem 5.1. Moreover, because it also implies boundedness of the cost under the proposed policies, it enables us to use Lemma 5.2. The proof is given in Section 5.2.2.

Proposition 5.4. *For every $t_0 > 0$, as $n \rightarrow \infty$,*

$$\mathbb{P}(\|\hat{X}_p^n\|_{t_0} \geq 2\hat{\Theta}^n) \rightarrow 0.$$

The above type of result is often referred to as state-space collapse (SSC), as it asserts that, asymptotically, all workload is kept in one buffer, a property crucially used in establishing the one-dimensional state space description of the limiting dynamics, as well as asymptotic optimality, because all workload is held in the “less expensive” class. It is proved in Section 5.2.3.

Proposition 5.5. *Consider a subsequence as in Lemma 5.2, where $(\hat{A}^n, \hat{S}^n, T^n) \Rightarrow (A, S, T)$. Then, along this sequence, one has $(\hat{A}^n, \hat{S}^n, T^n, \hat{F}^n, \hat{W}^n, \hat{L}^n) \Rightarrow (A, S, T, F, W, L)$, where F is given by (5.8), and $(W, L) = \Gamma(F)$. In particular, $(\hat{W}^n, \hat{L}^n)$ is a C -tight sequence, and the conclusions of Lemma 5.2 hold with $(W, L) = (Z', L')$.*

The proof appears in Sections 5.2.4 and 5.2.5.

Finally, the policies $\mathbf{P}$ and $\mathbf{T}_2$ that are proposed under the single-mode condition do not use the nonbasic activity of the active mode ξ^A . For the four policies employed under the dual-mode condition, we need the following control over the use of the nonbasic activities corresponding to ξ^L and ξ^H .

Proposition 5.6. *Consider the same subsequence as in Proposition 5.5 and cases 2(a)–(d) of Theorem 2.13. If (i^l, k^l) (resp. (i^h, k^h)) denotes the nonbasic activity in ξ^L (resp. ξ^H), then for any $t_0 > 0$ and $\varepsilon > 0$,*

$$\int_0^{t_0} \mathbb{1}_{\hat{W}_t^n \leq z^* - \varepsilon} dT_{i^l k^l}^n(t) \rightarrow 0, \quad \int_0^{t_0} \mathbb{1}_{\hat{W}_t^n \geq z^* + \varepsilon} dT_{i^h k^h}^n(t) \rightarrow 0, \quad (5.10)$$

in probability, as $n \rightarrow \infty$.

To explain the role of this proposition, recall that if $\xi \in S_{LP}$, then the condition $\xi_{ik} = 0$ for some activity (i, k) not only implies that ξ is one of the modes $\xi^{*,1}$ or $\xi^{*,2}$, but also identifies which one by Lemma 3.2. Because any limit T of T^n is given as $\int_0^\cdot \Xi_s ds$, $\Xi_t \in S_{LP}$, this proposition implies that when workload is either below or above the switching point, the resource allocation asymptotically follows the respective mode of operation ξ^L or ξ^H . This is

very close to stating that the pair (W, L) follows the SDE (2.27) and, indeed, is the basis for proving this fact. The proof of this proposition appears in Section 5.2.6.

5.1.3. Considerations for the Construction of Dual-Mode Policies. As noted above, two of the key results that we need to prove Theorem 2.13 are state-space collapse (Proposition 5.4) and a boundary property stating that there is asymptotically no idleness of either server when there is work in the system (Proposition 5.5). The proofs of these results differ by case/policy and, for dual-mode policies, rely on the rules used as well as the way that workload is sampled. Here, we provide a brief description of the reasons underlying the policy definitions that we have used.

- A difficulty arises in the proof of Proposition 5.4 in case 2(a), where the policy switches between two $\mathbf{P}$ rules. Only the dual-activity server in the current mode processes the HPC and the identity of the dual-activity server changes when switching modes. At each switching time, if the server processing the HPC in the new mode is busy with low-priority jobs and the service of the high-priority job at the other server ends, no server will process high-priority jobs for a time $O(n^{-1})$. In principle, this time could accumulate to let the number of high-priority jobs increase to a nonnegligible value. In order to prevent this, we designed switching between $\mathbf{P}$ rules to only occur at the time of service completion at the single-activity server. When switching, this server becomes dual activity and gives priority to the HPC (which is the single-activity class in $\mathbf{P}$). In case 2(a), there is always at least one server processing the HPC, regardless of switching between modes.

- Similarly, in order to prove Proposition 5.5 in case 2(b), we need to show that the number of HPC is not zero when there are low-priority jobs in the system. To make sure that the high-priority class does not receive too much service, switching between two $\mathbf{T}_2$ rules only occurs at the time of service completion at the single-activity server. When switching, the single-activity server becomes dual activity and now gives priority to the low-priority jobs if the number of HPC jobs is low. In case 2(b), as long as the number of HPC jobs is below the threshold there is at most one server working on HPC jobs, regardless of switching between modes.

- In case 2(c), it should be noted that in the lower mode, the system looks similar to case 1(a), so one could consider using a $\mathbf{P}$ rule. Doing so, however, would cause difficulty in proving Proposition 5.5 for this case. Using a $\mathbf{P}$ rule keeps the number of HPC jobs $O(1)$. At the moment of switching into the $\mathbf{T}_2$ rule, this can lead to the new single-activity server, which is dedicated to HPC, incurring idle time when there is work in the system. To avoid this situation, the $\mathbf{T}_1$ rule was used instead of $\mathbf{P}$, guaranteeing that, at a switching time, there will not be too few HPC jobs in the system. A similar argument applies to case 2(d).

5.1.4. Proof of Weak Convergence. Here, we prove Theorem 5.1 based on the four propositions.

Proof of Theorem 5.1. First, by Proposition 5.3, one has that $\limsup_n \int^n (T^n) < \infty$ for each one of the relevant policies T^n . As a result, the assumptions of Lemma 5.2 and Proposition 5.5 are valid. To summarize the conclusions from these results, fix a subsequence along which $(\hat{A}^n, \hat{S}^n, T^n) \Rightarrow (A, S, T)$. Then, there exists a tuple $(\{\mathcal{H}_t\}, F, W, L, \Xi, B)$ such that, along this sequence,

$$(\hat{A}^n, \hat{S}^n, \hat{F}^n, T^n, \hat{W}^n, \hat{L}^n) \Rightarrow (A, S, F, T, W, L),$$

where F is given in terms of A, S, T by (5.8) and W and L are given by $(W, L) = \Gamma(F)$. Moreover, $\mathfrak{S} = (\Omega, \mathcal{F}, \{\mathcal{H}_t\}, \mathbb{P}, B, \Xi, W, L)$ forms an admissible control for initial condition 0, and $T = \int_0^\infty \Xi_s ds$. In particular,

$$W = F + L = \int_0^\infty b(\Xi_s) ds + \int_0^\infty \sigma(\Xi_s) dB_s + L_t. \quad (5.11)$$

Next, by Proposition 5.4, $\hat{X}_p^n \rightarrow 0$ in probability. Also, by (2.11), $\hat{X}_q^n = \alpha_q \hat{W}^n - \alpha_p \hat{X}_p^n$, and, therefore, we now have, along the subsequence, $(\hat{A}^n, \hat{S}^n, \hat{F}^n, T^n, \hat{W}^n, \hat{L}^n, \hat{X}^n) \Rightarrow (A, S, F, T, W, L, X)$, where we denote

$$X_p = 0, \quad X_q = \alpha_q W. \quad (5.12)$$

Note that the control system $\mathfrak{S}$ thus constructed may depend on the subsequence. However, consider the following.

Claim. In case 1, (a)–(b) of Theorem 2.13, (W, L) is the RBM, and its boundary term given by (2.26), and $T_t = \xi^A t$ for all t . In cases 2(a)–(d) of Theorem 2.13, (W, L) is a weak solution to the SDE (2.27), and, moreover, $\Xi_t = \varphi^*(W_t)$ for a.e. t .

Suppose the above claim holds true. Then, the law of (W, L) is uniquely determined: in case 1 as an RBM; in case 2 as a weak solution to (2.27), for which weak uniqueness holds by Lemma 4.1.3. In particular, this law does not depend on the subsequence. Moreover, because by this claim and (5.12), the pair of processes (T, X) is

uniquely determined by W (away from a $\mathbb{P}$ -null set), it follows that the law of (T, W, L, X) does not depend on the subsequence. This yields the convergence $(T^n, \hat{W}^n, \hat{L}^n, \hat{X}^n) \Rightarrow (T, W, L, X)$ along the full sequence and completes the proof of the result. In what follows, the claim is proved.

Consider first the single-mode case, namely, cases 1(a)–(b) of Theorem 2.13. The policies employed are $\mathbf{P}$ and $\mathbf{T}_2$, and both do not use the nonbasic activity of the active mode ξ^A . In other words, if we denote this activity by (i^a, k^a) , then under these policies, $T_{i^a k^a}^n = 0$ for all n . As a consequence, the limit process T must satisfy $T_{i^a k^a}(t) = 0$ a.s.; hence, $\Xi_{i^a k^a}(t) = 0$ for a.e. t , a.s. By the uniqueness statement made in Lemma 3.2, whenever $\xi \in S_{LP}$ and $\xi_{i^a k^a} = 0$, one must have $\xi = \xi^A$. It follows that $\Xi_t = \xi^A$ for a.e. t , and, hence $T_t = \xi^A t$. As a consequence, (5.11) holds as $W_t = b(\xi^A)t + \sigma(\xi^A)B_t + L_t$. That is, the pair (W, L) satisfies (2.26). This proves the first part of the claim.

Next, consider the dual mode, namely, cases 2(a)–(d) of Theorem 2.13. First, we will show based on Proposition 5.6 that, for every $\varepsilon > 0$ and $t_0 > 0$,

$$\int_0^{t_0} \mathbb{1}_{\{W_t < z^* - \varepsilon\}} dT_{i^h k^h}(t) = 0, \quad \int_0^{t_0} \mathbb{1}_{\{W_t > z^* + \varepsilon\}} dT_{i^h k^h}(t) = 0, \quad (5.13)$$

holds a.s. Let g be a continuous function such that

$$\mathbb{1}_{w < z^* - \varepsilon} \leq g(w) \leq \mathbb{1}_{w < z^* - \frac{\varepsilon}{2}}, \quad w \in \mathbb{R}_+. \quad (5.14)$$

By the continuous mapping theorem, we have along the subsequence $(g(\hat{W}^n), T^n) \Rightarrow (g(W), T)$. In addition, T^n is continuous with finite variation over compacts. By theorem 2.2 of Kurtz and Protter (1991),

$$\int_0^\cdot g(\hat{W}_t^n) dT_{i^h k^h}^n(t) \Rightarrow \int_0^\cdot g(W_t) dT_{i^h k^h}(t). \quad (5.15)$$

By Proposition 5.6, $\int_0^{t_0} \mathbb{1}_{\{\hat{W}_t^n < z^* - \frac{\varepsilon}{2}\}} dT_{i^h k^h}^n(t) \rightarrow 0$ in probability. Hence, the left-hand side (LHS) in (5.15) converges to zero in probability. Thus, the RHS in (5.15) equals zero a.s., and, therefore, by the first inequality in (5.14), the first part of (5.13) is proved. The second part of (5.13) is proved analogously.

Next, clearly, $\Xi_{i^h k^h}(t) \mathbb{1}_{\{\Xi_t = \xi^L\}} = 0$. Hence, by (5.13), $\int_0^{t_0} \mathbb{1}_{\{W_t < z^* - \varepsilon\}} \mathbb{1}_{\{\Xi_t \neq \xi^L\}} \Xi_{i^h k^h}(t) dt = 0$. Arguing, again, by the uniqueness statement in Lemma 3.2, one has $\Xi_{i^h k^h}(t) > 0$ whenever $\Xi_t \neq \xi^L$. It follows that, a.s.,

$$\int_0^{t_0} \mathbb{1}_{\{W_t < z^* - \varepsilon\}} \mathbb{1}_{\{\Xi_t \neq \xi^L\}} dt = 0, \quad \int_0^{t_0} \mathbb{1}_{\{W_t > z^* + \varepsilon\}} \mathbb{1}_{\{\Xi_t \neq \xi^H\}} dt = 0, \quad (5.16)$$

where the second equality is proved analogously to the first one. Going back to (5.11), note that by (5.16), one has both $\int_0^{t_0} b(\Xi_s) \mathbb{1}_{\{W_t < z^* - \varepsilon\}} \mathbb{1}_{\{\Xi_t \neq \xi^L\}} ds = 0$ and $\int_0^{t_0} \sigma(\Xi_s) \mathbb{1}_{\{W_t < z^* - \varepsilon\}} \mathbb{1}_{\{\Xi_t \neq \xi^L\}} dB_s = 0$. A similar statement hold for $\{W_t > z^* + \varepsilon\}$ and ξ^H . Hence, by (5.11) and the Definition (2.28) of the functions b^* and σ^* , it follows that

$$\begin{aligned} W_t &= \int_0^t \mathbb{1}_{|W_s - z^*| \geq \varepsilon} b^*(W_s) ds + \int_0^t \mathbb{1}_{|W_s - z^*| \geq \varepsilon} \sigma^*(W_s) dB_s \\ &\quad + \int_0^t \mathbb{1}_{|W_s - z^*| < \varepsilon} b(\Xi_s) ds + \int_0^t \mathbb{1}_{|W_s - z^*| < \varepsilon} \sigma(\Xi_s) dB_s + L_t, \end{aligned}$$

for $t \leq t_0$. Because t_0 is arbitrary, this is true for all t . Hence, denoting $\delta_t^\varepsilon = \mathbb{1}_{|W_t - z^*| < \varepsilon}$ and using the boundedness of $b(\cdot)$,

$$U_t := \left| W_t - \int_0^t b^*(W_s) ds - \int_0^t \sigma^*(W_s) dB_s - L_t \right| \leq \gamma_t^\varepsilon + |\tilde{\gamma}_t^\varepsilon| + |\check{\gamma}_t^\varepsilon|,$$

where

$$\gamma_t^\varepsilon = c \int_0^t \delta_s^\varepsilon ds, \quad \tilde{\gamma}_t^\varepsilon = \int_0^t \delta_s^\varepsilon \sigma(\Xi_s) dB_s, \quad \check{\gamma}_t^\varepsilon = \int_0^t \delta_s^\varepsilon \sigma^*(W_s) dB_s.$$

Notice that U_t does not depend on ε , so if we manage to prove that the RHS converges to zero in probability as $\varepsilon \downarrow 0$, it follows that, for every t , $U_t = 0$ a.s. Hence, by continuity of this process, $U_t = 0$ for all t , a.s. That is, the processes (W, L, B) satisfy (2.27) a.s.

To this end, apply the occupation times formula (Revuz and Yor 1999, corollary VI.1.6), by which for any continuous semimartingale Y , one has, a.s.,

$$\int_0^t \mathbb{1}_{\{Y_s=0\}} d\langle Y, Y \rangle_s = \int_{-\infty}^{\infty} \mathbb{1}_{y=0} L_t^y(Y) dy,$$

with $L^y(Y)$ the local time of Y at y . Consider the above with $Y_t = W_t - z^*$. Clearly, the RHS in the above display is zero. Moreover, $\langle Y, Y \rangle_t = \int_0^t \sigma(\Xi_s)^2 ds$, and because σ is bounded away from zero, we obtain that $\int_0^t \mathbb{1}_{\{W_s=z^*\}} ds = 0$ a.s. For fixed ω , one has for all t , $\mathbb{1}_{\{|W_t-z^*|<\varepsilon\}} \rightarrow \mathbb{1}_{\{W_t=z^*\}}$ as $\varepsilon \downarrow 0$. Hence, by dominated convergence, $\gamma_t^\varepsilon \rightarrow 0$ a.s. as $\varepsilon \downarrow 0$.

Next, for the stochastic integral $\tilde{\gamma}^\varepsilon$, we have

$$\mathbb{E}(\tilde{\gamma}^\varepsilon)^2 = \mathbb{E} \int_0^t (\delta_s^\varepsilon \sigma(\Xi_s))^2 ds \leq c \mathbb{E} \int_0^t \delta_s^\varepsilon ds = c \mathbb{E}(\gamma_t^\varepsilon).$$

By local boundedness of γ^ε and its a.s. convergence to 0 as $\varepsilon \downarrow 0$, it follows that $\tilde{\gamma}_t^\varepsilon \rightarrow 0$ in probability. A similar argument holds for $\tilde{\gamma}_t^\varepsilon$. We conclude that (W, L, B) satisfies (2.27). Finally, by (5.16) and the fact $\int_0^t \mathbb{1}_{\{W_s=z^*\}} ds = 0$ a.s. just proved, one has $\Xi_t = \varphi^*(W_t)$ for a.e. t , a.s., which completes the proof of the claim, and also of the result. $\square$

5.1.5. Proof of Theorem 2.13 and Corollary 2.15. As a consequence of Theorem 5.1 and Proposition 5.3, we have the following.

Proof of Theorem 2.13. The weak convergence statements asserted in the theorem are already established in Theorem 5.1. For the AO results, we need to show that in each of the six cases $\hat{J}^n(T^n) \rightarrow V_0 = h_q \alpha_q V_{\text{WCP}}(0)$. Combining Theorem 5.1 with the identification of an optimal control for the WCP given in Lemma 4.1 shows that $\hat{X}^n \Rightarrow X$, where $X_p = 0$ and $X_q = \alpha_q W$, W is given by (2.26) or (2.27) in the respective cases, and, moreover,

$$V_0 = h_q \alpha_q \mathbb{E} \int_0^\infty e^{-\gamma t} W_t dt.$$

The convergence stated above implies

$$\hat{H}_t^n = h \cdot \hat{X}_t^n \Rightarrow h_q \alpha_q W_t. \quad (5.17)$$

By (2.6), $\hat{J}^n(T^n) = \mathbb{E} \int_0^\infty e^{-\gamma t} \hat{H}_t^n dt$. Hence, the convergence $\hat{J}^n(T^n) \rightarrow V_0$ will follow from (5.17) once uniform integrability is established. Arguing along the lines of Bell and Williams (2001, pp. 640–643), introduce the measure $dm = \gamma e^{-\gamma t} dt$ on $(\mathbb{R}_+, \mathcal{R}_+)$ and invoke the Skorohod representation theorem to obtain from (5.17) that $\hat{H}^n \rightarrow h_q \alpha_q W$ ($m \times \mathbb{P}$)-a.e. Accordingly, uniform integrability of H^n w.r.t. $m \times \mathbb{P}$ suffices to obtain $\hat{J}^n(T^n) \rightarrow V_0$. However, this is ensured by Proposition 5.3, for

$$\limsup_n \int_{[0, \infty) \times \Omega} (\hat{H}_t^n)^{1+\varepsilon_0} d(m \times \mathbb{P}) = \gamma \limsup_n \mathbb{E} \int_0^\infty e^{-\gamma t} (\hat{H}_t^n)^{1+\varepsilon_0} dt < \infty,$$

and the result is proved. $\square$

Proof of Corollary 2.15. As far as the proof of the Weak Convergence (2.29) and AO are concerned, there is no difference between the setting of Theorem 2.13(1), where the LP has multiple solutions but the single-mode condition holds, and the case where it has a single solution. Therefore, this result is a corollary of Theorem 2.13(1). $\square$

5.2. Proof of Propositions 5.3–5.6

In Section 5.2.1, we provide several useful estimates. Then, Propositions 5.3, 5.4, 5.5, and 5.6 are proved in Section 5.2.2, Section 5.2.3, Sections 5.2.4 and 5.2.5, and Section 5.2.6, respectively.

5.2.1. Auxiliary Lemmas. Two estimates on the rescaled primitives, $\hat{A}_i^n$ and $\hat{S}_{ik}^n$, are provided in (5.18) and Lemma 5.7, and a certain estimate on the maximum service duration is given in Lemma 5.9.

Because the assumptions on A_i^n and S_{ik}^n are similar, the estimates are stated for $\hat{A}_i^n$, but apply also for $\hat{S}_{ik}^n$. Recall that $\tilde{a}_i(l)$ are interarrival times of A_i and thus $a_i^n(l) := (\lambda_i^n)^{-1} \tilde{a}_i(l)$ are the interarrival times of A_i^n . Recall $E[\tilde{a}_i(1)^m] < \infty$ for a constant $m > 2$. The first estimate is Krichagina and Taksar (1992, theorem 4), which states that for any

$$2 \leq \kappa \leq \mathbf{m},$$

$$E[\|\hat{A}_i^n\|_t^\kappa] \leq c(1+t)^{\kappa/2}, \quad (5.18)$$

for a constant c that does not depend on n or t . The next useful estimate is as follows.

Lemma 5.7. *Let $v_1, v_2 \in (0, 1)$ be such that $v_1 \leq v_2 + \frac{1}{2}$ and assume that $h_0 := (\frac{\mathbf{m}}{2} - 1)v_1 - \mathbf{m}v_2 > 0$ (note, in particular, that for every $v_1 \in (0, \frac{1}{2}]$, there exists $v_2 > 0$ satisfying these conditions). Fix $c_1 > 0$. Then, for any $h < h_0$,*

$$P(w_{t_0}(\hat{A}_i^n, n^{-v_1}) \geq c_1 n^{-v_2}) \leq c n^{-h} (1+t_0)^{\mathbf{m}/2}, \quad n \in \mathbb{N}, t_0 \geq 1,$$

where $c = c(c_1, v_1, v_2, \mathbf{m})$ does not depend on n or t_0 .

The proof appears in Appendix A.

Remark 5.8. We sometimes use the balance equation in the following form, which follows from (5.4), namely,

$$\hat{X}_i^n(t) = \hat{f}_i^n(t) + n^{1/2} \int_0^t \left(\lambda_i - \sum_k \mu_{ik} \Xi_{ik}^n(s) \right) ds, \quad (5.19)$$

where

$$\hat{f}_i^n(t) := \hat{A}_i^n(t) - \sum_k \hat{S}_{ik}^n(T_{ik}^n(t)) + t \hat{\lambda}_i - \sum_k T_{ik}^n(t) \hat{\mu}_{ik}^n.$$

Note that $\hat{F}^n$ of (5.5) is given by $\hat{F}^n(t) = \sum_i \alpha_i^{-1} \hat{f}_i^n(t)$. Then, using the 1-Lipschitz property of the trajectories of T_{ik}^n , it is easy to see that both estimates above imply some estimates for $\hat{f}_i^n$. In particular, by (5.18), $\mathbb{E}[\|\hat{f}_i^n\|_t^\kappa] \leq c(1+t)^\kappa$. Moreover, under the assumptions of Lemma 5.7 and the additional assumption that $v_2 < v_1$, the conclusion of the lemma holds for $\hat{f}_i^n$ and $\hat{F}^n$.

Next, we give an estimate on the maximal service duration and interarrival time up to a given time. The *time in service by t* of a given job is defined as the time that the job has spent in service up to time t . Let $\text{TIS}(n, i, k, l, t)$ denote the time in service by t of the l th job in activity (i, k) . If service to job l has completed by time t , then, clearly, $\text{TIS}(n, i, k, l, t) = u_{ik}^n(l)$, but if it is still in service, $\text{TIS}(n, i, k, l, t) < u_{ik}^n(l)$. Of course, $\text{TIS}(n, i, k, l, t) = 0$ for jobs for which service has not started by t . For $t > 0$ and a real-valued path φ , denote

$$\Lambda(\varphi, t) = \sup\{t_2 - t_1 : 0 \leq t_1 \leq t_2 \leq t, \varphi_{t_2} = \varphi_{t_1}\}.$$

Then, for activity (i, k) , the maximal time in service by time t , namely, $\sup_l \text{TIS}(n, i, k, l, t)$, is bounded above by $\Lambda(S_{ik}^n, t)$. We will need an upper bound on the service time as well as the interarrival times, and to this end define

$$e_{\max}^n(t) = \max_{i,k} \Lambda(S_{ik}^n, t) \vee \max_i \Lambda(A_i^n, t). \quad (5.20)$$

This process also bounds from above all service durations completed by time t .

Lemma 5.9. *One has $\mathbb{E}[(e_{\max}^n)^2] \leq c n^{-1}(1+t)$ for a constant c that does not depend on n or t . Moreover, for any $t < \infty$ and $c_1 > 0$, $\mathbb{P}(e_{\max}^n(t) \geq c_1 n^{a-1}) \rightarrow 0$, as $n \rightarrow +\infty$.*

The proof appears in Appendix A.

5.2.2. Uniform Integrability. Here, we prove Proposition 5.3. We sometimes need to refer to the variable keeping track of the current mode, which in the various policies is defined in slightly different ways according to different sampling times. Denote

$$\text{MODE}^n(t) = \begin{cases} \xi^L & \text{if the current mode is lower workload,} \\ \xi^H & \text{if the current mode is higher workload.} \end{cases}$$

Proof of Proposition 5.3. The proof has three parts; part 1 appears below, while parts 2 and 3 are deferred to Appendix A. Fix any one of the sequences of policies $T^n \in \mathcal{A}^n$ for which we attempt to prove AO.

Part 1. This part is concerned with the case where, whenever the workload in the system is sufficiently large, policy **P** is active. This covers case 1(a) of Theorem 2.13, where the system has a single mode and the fixed priority policy **P** is applied, as well as case 2(a), where the dual-mode policy **PP** is applied. In this case, we will prove

that the statement of the lemma holds with $\varepsilon_0 = 1$, and specifically that, under the given sequence, $\mathbb{E}[(\hat{H}_t^n)^2]$ is bounded by a polynomial in t for all n . By (2.11), this is equivalent to the same property holding for $\mathbb{E}[(\hat{W}_t^n)^2]$.

To prove the result in this case, assume that in the single (resp., dual)-mode case, the active mode ξ^A (resp., the high-workload mode ξ^H) is in canonical form. Thus, provided that the workload in the system exceeds z^* (when applicable), class 1 (resp., 2) is the dual (single)-activity class, and server 1 (resp., 2) is the single (dual)-activity server. In addition, $p = 2$: class 2 is the HPC. First, we provide a bound on the second moment of $\hat{X}_2^n$. In the dual-mode case, fix K large enough so that $\hat{X}_2^n(t) \geq K$ implies $\hat{W}_t^n \geq z^* + 1$; in the single-mode case let $K = 1$. Given $t > 0$, consider the event $\hat{X}_2^n(t) > K$. Let $\tau = \tau_n(t)$ be defined by

$$\tau = \sup\{s \in [0, t] : \hat{X}_2^n(s) \leq K\}.$$

Because the system starts empty, $0 \leq \tau \leq t$, and because the jumps of the normalized queue length are of size $n^{-1/2}$, $\hat{X}_2^n(\tau) \leq K + 1$. Thus,

$$\hat{X}_2^n(t) \leq \hat{X}_2^n(t) - \hat{X}_2^n(\tau) + K + 1.$$

By (5.3), denoting $C^n(t) = \sum_i \|\hat{A}_i^n\|_t + \sum_{ik} \|\hat{S}_{ik}^n\|_t$,

$$\hat{X}_2^n(t) - \hat{X}_2^n(\tau) \leq 2C^n(t) + n^{-1/2}\lambda_2^n(t - \tau) - n^{-1/2}\mu_{22}^n(T_{22}^n(t) - T_{22}^n(\tau)).$$

Because $\hat{X}_2^n \geq K$ in $[\tau, t]$, the priority policy corresponding to a mode given in canonical form (either ξ^A or ξ^H) is in force after an initial time before the current mode updates and there are class 2 jobs to serve. We can bound the time until server 2 prioritizes class 2 and starts serving it at full rate in $[\tau, t]$ by $e_{\max}^n(t)$. If there currently is a class 2 job being served by server 2, server 2 only serves class 2 jobs for the whole period and

$$T_{22}^n(t) - T_{22}^n(\tau) = t - \tau.$$

If there currently is a class 1 job being served by server 2, because service is nonpreemptive, server 2 has to finish this job. If at that time, the current mode is ξ^H (resp. ξ^A in the single-mode case), server 2 prioritizes class 2 from that point on. If at that time, the current mode is ξ^L , the workload is then sampled because a service finished at the single activity server. The current mode changes to ξ^H and stays until t . In both cases, we get

$$T_{22}^n(t) - T_{22}^n(\tau) \geq t - \tau - e_{\max}^n.$$

Also, in the case under consideration, one has $\mu_{22} > \lambda_2$. Hence, for all sufficiently large n , $\mu_{22}^n > \lambda_2^n$. Moreover, $c_1 := \sup_n n^{-1}\mu_{22}^n < \infty$. Hence,

$$\hat{X}_2^n(t) \leq 2C^n(t) + c_1 n^{1/2} e_{\max}^n(t) + K + 1.$$

The above inequality holds also on the complementary event, namely, when $\hat{X}_2^n(t) \leq K$. Thus, using (5.18) and Lemma 5.9, we obtain

$$\mathbb{E}[\|\hat{X}_2^n\|_t^2] \leq c(1 + t). \quad (5.21)$$

In the next step, we bound $\hat{W}_t^n$. Let $\tilde{K}$ be a constant that is sufficiently large to ensure that $\hat{X}_1^n(t) \geq \tilde{K}$ implies $\hat{W}_t^n \geq z^* + 1$ (where, again, $\tilde{K} = 1$ in the single-mode case). Given $t > 0$, consider the event $\hat{X}_1^n(t) > \tilde{K}$ and let $\sigma = \sigma_n(t)$ be

$$\sigma = \sup\{s \in [0, t] : \hat{X}_1^n(s) \leq \tilde{K}\}.$$

Then, clearly, $\sigma \in [0, t]$. Because during $[\sigma, t]$, there are at least two jobs of class 1 in the system, both servers are never idle during $[\sigma, t]$. Because one has $\hat{X}_1^n(\sigma) \leq \tilde{K} + n^{-1/2}$, it follows that $\hat{W}_\sigma^n \leq c(1 + \tilde{K} + \hat{X}_2^n(\sigma))$. Hence,

$$\hat{W}_t^n \leq \hat{W}_t^n - \hat{W}_\sigma^n + c(1 + \tilde{K} + \|\hat{X}_2^n\|_t).$$

By (5.6) and the nonidling of both servers during $[\sigma, t]$, by which $\hat{L}^n$ remains flat during this interval, we have $\hat{W}_t^n - \hat{W}_\sigma^n = \hat{F}_t^n - \hat{F}_\sigma^n$. Thus, by (5.5), for a constant c (that may depend on $\tilde{K}$),

$$\hat{W}_t^n \leq c(C^n(t) + t + 1 + \|\hat{X}_2^n\|_t).$$

Clearly, the above bound is valid also in the complementary event, $\hat{X}_1^n(t) \leq \tilde{K}$. We can therefore apply (5.18) and (5.21), to obtain $\mathbb{E}[(\hat{W}_t^n)^2] \leq c(1 + t)$ for some constant c , for all n and t . Consequently, the same holds for the

second moment of $\hat{H}_t^n$, and the result follows. This completes part 1 of the proof. The remaining parts appear in Appendix A. $\square$

5.2.3. State Space Collapse. We now prove Proposition 5.4. Assume that either the active mode ξ^A or the lower workload mode ξ^L (whichever is applicable) is in canonical form. Let

$$\tau = \tau_c^n = \inf\{t \geq 0 : \hat{X}_p^n(t) \geq 2\hat{\Theta}^n\}.$$

This random time is used outside this proof with the notation τ_c^n , but in this proof, the shorter notation τ is used. Then,

$$\mathbb{P}\left(\sup_{t \leq t_0} \hat{X}_p^n(t) \geq 2\hat{\Theta}^n\right) \leq \mathbb{P}(\tau \leq t_0).$$

We will prove the lemma by showing that the RHS above converges to zero as $n \rightarrow \infty$. On the event $\{\tau \leq t_0\}$, define

$$\sigma = \sigma^n = \sup\left\{t \leq \tau : \hat{X}_p^n(t) \leq \frac{3\hat{\Theta}^n}{2}\right\}.$$

The proof relies on the fact that, under our policies, when the number of HPC jobs is above Θ^n , it is served at a rate that is enough to deplete the queue in all cases. The first step toward this goal is the following.

Lemma 5.10. *There exist constants $c_1, c_2 > 0$ such that, on $\{\tau \leq t_0\}$,*

$$\int_{\sigma}^{\tau} \left(\lambda_p - \sum_k \mu_{pk} \Xi_{pk}^n(t) \right) dt \leq -c_1(\tau - \sigma) + c_2 e_{\max}^n, \quad (5.22)$$

with $e_{\max}^n$ defined in (5.20).

Proof. First, by definition of σ , and τ , and the fact that jumps of $\hat{X}^n$ are of size $n^{-1/2} < \hat{\Theta}^n/2$ for large n , we obtain

$$\inf_{t \in [\sigma, \tau]} \hat{X}_p^n(t) \geq \hat{\Theta}^n.$$

We address here one case only; the proof under the remaining cases is deferred to Appendix A.

• Case 1(a): In this case, we use the **P** policy, which is single mode. Because the active mode is in canonical form and $i_1(\xi^A) = p$, $p = 2$, and server 2 prioritizes class 2. It is possible that server 2 is busy with the “wrong” class of job at time σ , but as soon as that job finishes, server 2 will only serve the class 2 jobs in the system:

$$\int_{\sigma}^{\tau} \Xi_{22}^n(t) dt \geq \tau - \sigma - e_{\max}^n.$$

Thus,

$$\int_{\sigma}^{\tau} \left(\lambda_2 - \sum_k \mu_{2k} \Xi_{2k}^n(t) \right) dt \leq (\tau - \sigma)(\lambda_2 - \mu_{22}) + \mu_{22} e_{\max}^n.$$

To show that $\lambda_2 - \mu_{22} < 0$, note that since the active mode is in canonical form, $\frac{\lambda_1}{\alpha_1} > \beta_1$, which implies $\frac{\lambda_2}{\alpha_2} < \beta_2$ by (3.1).

The remaining cases appear in Appendix A. $\square$

Now that Lemma 5.10 has been proved in all cases, the proof of Proposition 5.4 does not need to differentiate between them.

Proof of Proposition 5.4. Let $v_2 = 1/2 - \bar{a} \leq 1/4$ and $v_1 \in (v_2, 1/2)$. Recall that $\hat{\Theta}^n = n^{-1/2} \lceil n^{\bar{a}} \rceil$, as defined in (2.32), so that $\hat{\Theta}^n = n^{-1/2} \lceil n^{1/2-v_2} \rceil \geq n^{-v_2}$. Notice that, as required in Lemma 5.7, $v_1 \leq v_2 + \frac{1}{2}$. Let us introduce the event

$$\Omega_1 = \left\{ \tau \leq t_0, \hat{X}_p^n(\tau) - \hat{X}_p^n(\sigma) \geq \frac{\hat{\Theta}^n}{4}, \inf_{t \in [\sigma, \tau]} \hat{X}_p^n(t) \geq \hat{\Theta}^n \right\}.$$

For n large enough, one has $\hat{X}_p^n(\sigma) < \frac{7}{4}\hat{\Theta}^n$. Consequently, $\{\tau \leq t_0\} = \Omega_1$. We obtain

$$\mathbb{P}(\tau \leq t_0) = \mathbb{P}(\Omega_1) = \mathbb{P}(\Omega_1 \cap \{\tau - \sigma \leq n^{-\nu_1}\}) + \mathbb{P}(\Omega_1 \cap \{\tau - \sigma > n^{-\nu_1}\}). \quad (5.23)$$

Picking up on (5.19), by Lemma 5.10

$$\begin{aligned} \hat{X}_p^n(\tau) - \hat{X}_p^n(\sigma) &= \hat{f}_p^n(\tau) - \hat{f}_p^n(\sigma) + \sqrt{n} \int_{\sigma}^{\tau} \left(\lambda_p - \sum_k \mu_{pk} \Xi_{pk}^n(t) \right) dt \\ &\leq \hat{f}_p^n(\tau) - \hat{f}_p^n(\sigma) - c\sqrt{n}(\tau - \sigma) + c\sqrt{n}e_{\max}^n. \end{aligned}$$

Notice that on the event $\{\tau - \sigma \leq n^{-\nu_1}, \tau \leq t_0\}$,

$$\hat{X}_p^n(\tau) - \hat{X}_p^n(\sigma) \leq w_{t_0}(\hat{f}_p^n, n^{-\nu_1}) + c\sqrt{n}e_{\max}^n.$$

Thus,

$$\begin{aligned} \mathbb{P}(\Omega_1 \cap \{\tau - \sigma \leq n^{-\nu_1}\}) &\leq \mathbb{P}\left(w_{t_0}(\hat{f}_p^n, n^{-\nu_1}) + e_{\max}^n \sqrt{n} \geq \frac{n^{-\nu_2}}{2}\right) \\ &\leq \mathbb{P}\left(w_{t_0}(\hat{f}_p^n, n^{-\nu_1}) \geq \frac{n^{-\nu_2}}{4}\right) + \mathbb{P}\left(ce_{\max}^n \sqrt{n} \geq \frac{n^{-\nu_2}}{4}\right). \end{aligned}$$

Recall that $\nu_1 \in (\nu_2, 1/2)$, so $n^{-\nu_1} < n^{-\nu_2}$. By Lemma 5.7 and Remark 5.8, the first term converges to zero. By Lemma 5.9 and $\nu_2 = 1/2 - \bar{a}$,

$$\mathbb{P}\left(e_{\max}^n \geq \frac{c}{4}n^{-\nu_2-1/2}\right) = \mathbb{P}\left(e_{\max}^n \geq \frac{c}{4}n^{\bar{a}-1}\right) \rightarrow 0.$$

For the second term in (5.23), notice that on $\{\tau - \sigma > n^{-\nu_1}, \tau \leq t_0\}$,

$$\hat{X}_p^n(\tau) - \hat{X}_p^n(\sigma) \leq 2\|\hat{f}_p^n(t)\|_{t_0} + e_{\max}^n \sqrt{n} - c\sqrt{nn}^{-\nu_1},$$

and

$$\hat{X}_p^n(\tau) - \hat{X}_p^n(\sigma) \geq \frac{\hat{\Theta}^n}{2} - n^{-1/2}.$$

Hence,

$$\mathbb{P}(\Omega_1 \cap \{\tau - \sigma > n^{-\nu_1}\}) \leq \mathbb{P}\left(2\|\hat{f}_p^n\|_{t_0} + e_{\max}^n \sqrt{n} \geq c\sqrt{nn}^{-\nu_1} + \frac{\hat{\Theta}^n}{2} - n^{-1/2}\right).$$

By Lemma 5.9, $e_{\max}^n$ is smaller than $n^{\bar{a}-1} = o(n^{-1/2})$. By Remark 5.8, $2\|\hat{f}_p^n\|_{t_0}$ is a tight sequence of random variables (RVs) (for t_0 fixed). Because $\nu_1 < 1/2$, $\sqrt{nn}^{-\nu_1} \rightarrow +\infty$. The claim follows. $\square$

5.2.4. Boundary Behavior. The goal of this section and the following one is to prove Proposition 5.5. The key is the following lemma, which states, roughly speaking, that $\hat{L}^n$ is approximately the boundary term corresponding to $\hat{F}^n$.

Let $c_3 = \frac{3}{\alpha_1 \wedge \alpha_2}$.

Lemma 5.11. Fix $t_0 > 0$. With

$$\bar{R}_t^n = \int_0^t \mathbf{1}_{\hat{W}_s^n \geq c_3 \hat{\Theta}^n} d\hat{L}_s^n, \quad (5.24)$$

$\bar{R}_{t_0}^n \rightarrow 0$ in probability as $n \rightarrow \infty$.

This lemma is proved in the next subsection. Let us show how Proposition 5.5 now follows.

Proof of Proposition 5.5. In addition to $\bar{R}^n$ just introduced, we shall need the following definitions:

$$\begin{aligned} Z_t^n &= \max(\hat{W}_t^n - c_3 \hat{\Theta}^n, 0), \\ \tilde{R}_t^n &= \int_0^t \mathbb{1}_{\hat{W}_s^n < c_3 \hat{\Theta}^n} d\hat{L}_s^n, \\ \Delta_t^n &= Z_t^n - \hat{W}_t^n. \end{aligned}$$

By (5.6) and definition of the new processes, one has

$$Z_t^n = \hat{F}_t^n + \Delta_t^n + \bar{R}_t^n + \tilde{R}_t^n.$$

Consider the pair $(Z^n, \tilde{R}^n)$. The first component is nonnegative. The second is nonnegative, nondecreasing, and, moreover,

$$\int_{[0, \infty)} Z_t^n d\tilde{R}_t^n = \int_{[0, \infty)} \max(\hat{W}_t^n - c_3 \hat{\Theta}^n, 0) \mathbb{1}_{\hat{W}_t^n < c_3 \hat{\Theta}^n} d\hat{L}_t^n = 0.$$

Hence, by Skorohod's Lemma, $(Z^n, \tilde{R}^n) = \Gamma(\hat{F}^n + \Delta^n + \bar{R}^n)$. Hence,

$$\hat{W}^n = \Gamma_1(\hat{F}^n + \Delta^n + \bar{R}^n) - \Delta^n, \quad \hat{L}^n = \Gamma_2(\hat{F}^n + \Delta^n + \bar{R}^n) + \bar{R}^n. \quad (5.25)$$

Recall that we are considering a subsequence along which (according to the assumptions and to Lemma 5.2), $(\hat{A}^n, \hat{S}^n, T^n, \hat{F}^n) \Rightarrow (A, S, T, F)$, with F as in (5.8). Moreover, by the definition of Z^n and Δ^n , we have $0 \leq -\Delta_t^n \leq c_3 \hat{\Theta}^n$, whereas by Lemma 5.11, $\bar{R}^n \rightarrow 0$ in probability. By the continuity of the map Γ , we therefore conclude from (5.25) that, on the same subsequence,

$$(\hat{A}^n, \hat{S}^n, T^n, \hat{F}^n, \hat{W}^n, \hat{L}^n) \Rightarrow (A, S, T, F, W, L),$$

where $(W, L) = \Gamma(F)$. This completes the proof. $\square$

5.2.5. Proof of Lemma 5.11. We first explain how the proof changes between the cases.

- Under the **P** policy, both servers can process the low priority jobs. This means that idling can only occur if the low-priority class has few jobs, and, in that case, the total number of jobs is also low by Proposition 5.4.
- Under the **T₂** rule, no server can incur idleness when the high-priority class is above the threshold. In addition, the high-priority class takes a small time to reach above Θ^n and does not empty below two jobs after that time with high probability unless the total workload is close to zero (Lemma 5.13). This means that neither server will idle except when there are almost no jobs in the system.
- Under the **PP** policy, Proposition 5.4 is still enough to prove Lemma 5.11 in the same way as in the single-mode case **P**.
- Under the **T₂T₂** policy, Lemma 5.12 and 5.13 hold. Because of that, Lemma 5.11 holds for the same reason as in the nonswitching case **T₂**.
- When using a **P** rule, the high-priority class could become zero with a lot of LPC jobs in the system, and switching to the **T₂** rule could lead to idleness, even though there are a lot of low-priority jobs (approximately $\alpha_q z^*$). Thus, we introduce the **T₁** rule in place of the **P** rule in this case. This ensures that the number of high-priority jobs does not decrease too much during the corresponding period.

In some cases, some idleness can occur when the high-priority class starts with too few jobs, but in those cases, it takes a small time to leave such states.

In cases 1(a) and 2(a), Lemma 5.11 is a direct consequence of the state space collapse. Let us introduce two random times and a lemma that we will use in cases 1(b) and 2(b)–(d). Fix $t_0 > 0$. Let

$$\begin{aligned} \rho^n &= \inf\{t \geq 0 : \hat{X}_p^n(t) \geq \hat{\Theta}^n\} \wedge t_0, \\ \tau_r^n &= \inf\{t > \rho^n : \hat{X}_p^n(t) = 2n^{-1/2}, \text{ and } \hat{X}_q^n(t) \geq \hat{\Theta}^n\}. \end{aligned}$$

Lemma 5.12. *If any of the following conditions hold, we have*

$$\bar{R}_{\rho^n}^n \rightarrow 0 \text{ in probability.} \quad (5.26)$$

- Case 1(b), (2.20) and $i_2(\xi^A) = p$.
- Case 2(b), (2.21) and $i_2(\xi^L) = i_2(\xi^H) = p$.
- Cases 2(c) and 2(d), (2.21) and $i_1(\xi^L) = i_2(\xi^H)$.

Lemma 5.13. *Under the same assumptions as the previous lemma,*

$$\mathbb{P}(\tau_r^n \leq t_0) \rightarrow 0. \quad (5.27)$$

The proofs of Lemmas 5.12 and 5.13 appear in Appendix A.

Remark 5.14. The above two lemmas do not address cases 1(a) and 2(a). The reason for this is that in these cases, we can directly prove Lemma 5.11 when using a **P** or **PP** policy.

Proof of Lemma 5.11. Here, we only treat one case, deferring the remaining cases to Appendix A.

- Case 1(a), P policy:

In this case, $p = 2$, and no server can be idle when $X_1 \geq 2$. Thus,

$$\int_0^{t_0} \mathbb{1}_{\hat{X}_1^n(t) \geq 2n^{-1/2}} d\hat{L}_t^n = 0. \quad (5.28)$$

For any $\delta > 0$,

$$\begin{aligned} \mathbb{P}(\bar{R}_{t_0}^n \geq \delta) &\leq \mathbb{P}\left(\int_0^{t_0} (\mathbb{1}_{\hat{X}_2^n(t) \geq 2\hat{\Theta}^n} + \mathbb{1}_{\hat{X}_1^n(t) \geq 2n^{-1/2}}) d\hat{L}_t^n \geq \delta\right) \\ &= \mathbb{P}\left(\int_0^{t_0} \mathbb{1}_{\hat{X}_2^n(t) \geq 2\hat{\Theta}^n} d\hat{L}_t^n \geq \delta\right) \\ &\leq \mathbb{P}\left(\mathbb{1}_{\sup_{t \leq t_0} \hat{X}_2^n(t) \geq 2\hat{\Theta}^n} \hat{L}^n(t_0) \geq \delta\right) \\ &\leq \mathbb{P}\left(\sup_{t \leq t_0} \hat{X}_2^n(t) \geq 2\hat{\Theta}^n\right) \\ &= \mathbb{P}(\tau_c^n \leq t_0). \end{aligned}$$

By Proposition 5.4,

$$\mathbb{P}(\tau_c^n \leq t_0) \rightarrow 0. \quad (5.29)$$

The treatment of the remaining cases appears in Appendix A. $\square$

5.2.6. Fast Switching

Proof of Proposition 5.6. Fix t_0 and $\varepsilon > 0$. Assume that ξ^L is in canonical form. Then, $(i^l, k^l) = (2, 1)$. Let

$$\begin{aligned} \tau_f &= \inf\left\{t \geq 0 : \int_0^t \mathbb{1}_{\hat{W}_t^n \leq z^* - \varepsilon} dT_{21}^n(t) > 0\right\}, \\ t_{\min} &= \sup\{t \leq \tau_f : \hat{W}_t^n \geq z^*\}, \\ t_{\max} &= \inf\{t \geq t_{\min} : \hat{W}_t^n \leq z^* - \varepsilon\}, \\ \tau_1 &= \inf\{t \geq t_{\min} : \text{MODE}^n(t) = \xi^L\}, \end{aligned}$$

where the dependence on n is suppressed. The first statement of the lemma will be proved once we show $\mathbb{P}(\tau_f \leq t_0) \rightarrow 0$. We omit the proof of the second statement, which is similar. To this end, let

$$\kappa^n = \inf\{s \geq 0 \text{ s.t. } \exists t \leq t_0, \hat{W}_{t-s}^n \geq z^*, \hat{W}_t^n \leq z^* - \varepsilon\}.$$

If $\sup_{t \leq t_0} \hat{W}_t^n \leq z^*$, the current mode never changes, so the single-activity server is dedicated to only one class for the whole period and the nonbasic activity is never used. Thus,

$$\{\tau_f \leq t_0\} \subset \{t_{\max} \leq t_0\}.$$

Note that we have used the fact that, for all of our policies, jobs are never routed to a nonbasic activity of the current mode.

The remainder of the argument is based on the following fact, which must be argued separately for each case. This is concerned with the difference $\tau_1 - t_{\min}$, which, although case-dependent, can in all cases be shown to

satisfy

$$\lim_{n \rightarrow +\infty} \mathbb{P}(\tau_1 - t_{\min} > e_{\max}^n) = 0. \quad (5.30)$$

By Proposition 5.5, $\hat{W}^n$ is C-tight. Hence, for any constant $c > 0$, $\mathbb{P}(c\kappa^n \leq n^{\bar{a}-1}) \rightarrow 0$. On the other hand, by Lemma 5.9, $\mathbb{P}(e_{\max}^n \geq n^{\bar{a}-1}) \rightarrow 0$. Thus,

$$\mathbb{P}(e_{\max}^n \geq c\kappa^n) \rightarrow 0. \quad (5.31)$$

In order for $\int_0^{t_0} \mathbb{1}_{W_t^n \leq z^* - \varepsilon} dT_{21}^n(t)$ to become positive, $t_{\max}$ must be smaller than t_0 and

$$t_{\max} - \tau_1 < e_{\max}^n.$$

It is not possible for $t_{\max} - \tau_1 \geq e_{\max}^n$ to occur on $\{\tau_f \leq t_0\}$. If that were the case, server 1 would necessarily finish service of the job it was in the process of serving at time τ_1 before the workload has time to reach $z^* - \varepsilon$. When $\text{MODE}^n(t) = \xi^L$ in canonical form, server 1 can only take new class 1 jobs, regardless of the rule. Even if class 1 has no job in the queue, the nonbasic activity is not used after the possible residual job that was in service at time τ_1 . In addition, by definition of $t_{\min}$, this is the last time the mode switches from upper to lower workload before τ_f . This would prevent $\int_0^{\tau_f} \mathbb{1}_{W_t^n \leq z^* - \varepsilon} dT_{21}^n(t)$ from becoming positive and is also the reason why $t_{\max}$ needs to be smaller than t_0 for τ_f to be smaller than t_0 .

By definition of $t_{\max}$ and κ^n ,

$$\mathbb{P}(t_{\max} \leq t_0, t_{\max} - t_{\min} < \kappa^n) = 0.$$

We next show that

$$\lim_{n \rightarrow +\infty} \mathbb{P}(t_{\max} \leq t_0, \tau_1 \geq t_{\max}) = 0. \quad (5.32)$$

We have

$$\begin{aligned} \mathbb{P}(t_{\max} \leq t_0, \tau_1 \geq t_{\max}) &= \mathbb{P}(t_{\max} \leq t_0, \tau_1 \geq t_{\max}, t_{\max} - t_{\min} \geq \kappa^n) \\ &\leq \mathbb{P}(\tau_1 \geq t_{\min} + \kappa^n) \\ &\leq \mathbb{P}(\tau_1 \geq t_{\min} + \kappa^n, e_{\max}^n \leq \kappa^n) + \mathbb{P}(\tau_1 \geq t_{\min} + \kappa^n, e_{\max}^n > \kappa^n) \\ &\leq \mathbb{P}(\tau_1 \geq t_{\min} + e_{\max}^n) + \mathbb{P}(e_{\max}^n > \kappa^n). \end{aligned}$$

Both terms converge to zero, the first by (5.30) and the second by (5.31).

We can now prove the lemma based on (5.32), (5.30), and (5.31). We have

$$\begin{aligned} \mathbb{P}(\tau_f \leq t_0) &= \mathbb{P}(\tau_f \leq t_0, t_{\max} \leq t_0) \\ &= \mathbb{P}(\tau_f \leq t_0, t_{\max} \leq t_0, \tau_1 \geq t_{\max}) + \mathbb{P}(\tau_f \leq t_0, t_{\max} \leq t_0, \tau_1 \leq t_{\max}, t_{\max} - \tau_1 < e_{\max}^n) \\ &\leq \mathbb{P}(t_{\max} \leq t_0, \tau_1 \geq t_{\max}) + \mathbb{P}(t_{\max} \leq t_0, t_{\max} - t_{\min} \leq 2e_{\max}^n) + \mathbb{P}(\tau_1 - t_{\min} > e_{\max}^n) \\ &\leq \mathbb{P}(t_{\max} \leq t_0, \tau_1 \geq t_{\max}) + \mathbb{P}(\kappa^n \leq 2e_{\max}^n) + \mathbb{P}(t_{\max} \leq t_0, t_{\max} - t_{\min} < \kappa^n) \\ &\quad + \mathbb{P}(\tau_1 - t_{\min} > e_{\max}^n). \end{aligned}$$

The first term goes to zero by (5.32), the second by (5.31), the third is zero, and the fourth goes to zero by (5.30). This completes the proof.

It remains to prove (5.30). As noted above, the proof differs by case.

Case 2(a): The current mode changes to ξ^L after the first service completion of a job at the single-activity server for ξ^H (which is server 2) at or after $t_{\min}$. Note that $t_{\min}$ must correspond to a service completion. If $t_{\min}$ corresponds to a service completion at server 2, the $\tau_1 = t_{\min}$. If $t_{\min}$ corresponds to a service completion at server 1, then the mode will not change, and τ_1 will correspond to the next service completion at server 2. This will occur before $t_{\min} + e_{\max}^n$ if server 2 is busy (serving class 1) at $t_{\min}$. Note that, on $\{\tau_c^n \geq t_0\}$, $\hat{X}_2^n(t_{\min}) < 2\hat{\Theta}^n$, so that

$$X_1^n(t_{\min}) \geq \alpha_1(W_{t_{\min}}^n - 2\alpha_2^{-1}\Theta^n) \geq \alpha_1(n^{1/2}z^* - 1 - 2\alpha_2^{-1}\Theta^n) \geq 2,$$

so server 2 is busy at $t_{\min}$. Thus,

$$\mathbb{P}(\tau_1 - t_{\min} > e_{\max}^n) \leq \mathbb{P}(\tau_c^n \leq t_0),$$

which converges to zero by Proposition 5.4. This proves (5.30).

Case 2(b): The current mode changes after the first service completion of a job at the single-activity server (which is server 2) at or after $t_{\min}$. Server 2 is dedicated to HPC jobs. Under $\{\tau_r^n \geq t_0, \rho^n \leq t_{\min}\}$, there are at least two HPC jobs in the system at time $t_{\min}$, so the single-activity server cannot be idling at that time. Thus,

$$\mathbb{P}(\tau_1 - t_{\min} > e_{\max}^n) \leq \mathbb{P}(\tau_r^n \leq t_0) + \mathbb{P}(\rho^n > t_{\min}).$$

By Lemma 5.13, $\mathbb{P}(\tau_r^n \leq t_0)$ converges to zero. In addition,

$$\begin{aligned} \mathbb{P}(\rho^n > t_{\min}) &= \mathbb{P}(\rho^n > t_{\min}, t_{\min} > \tilde{\tau}) + \mathbb{P}(\rho^n > t_{\min}, t_{\min} \leq \tilde{\tau}) \\ &\leq \mathbb{P}(\rho^n > \tilde{\tau}) + \mathbb{P}(t_{\min} \leq \tilde{\tau}) \\ &\leq \mathbb{P}(\rho^n > \tilde{\tau}, \tau_c^n \geq t_0) + \mathbb{P}(\rho^n > \tilde{\tau}, \tau_c^n < t_0) + \mathbb{P}(\hat{X}_2^n(t_{\min}) < \alpha_2 z^*/2) \\ &\leq \mathbb{P}(\rho^n > \tilde{\tau}, \tau_c^n \geq t_0) + 2\mathbb{P}(\tau_c^n < t_0). \end{aligned}$$

where $\tilde{\tau}$ is defined in the proof of Lemma 5.12. Both terms converge to zero, the first by Lemma 5.13, and the second by Proposition 5.4. This proves (5.30).

Cases 2(c) and 2(d): The current mode changes after the first service completion or arrival of a low- or high-priority job after $t_{\min}$ so under $\{t_{\max} \leq t_0\}$,

$$\tau_1 - t_{\min} \leq e_{\max}^n,$$

Thus, (5.30) follows from Lemma 5.9. $\square$

Acknowledgements

The authors are grateful to Cristina Costantini for providing us with the proof of Lemma 4.1, part 3, and to two anonymous referees, as well as the anonymous associate editor, for their thoughtful suggestions that helped improve the exposition considerably.

Appendix A. Proofs of Lemmas

Proof of Lemma 5.7. In this proof, $\mathbf{m}$ is written as m . Note first that it suffices to prove

$$\mathbb{P}(w_{t_0}(\hat{A}^n, n^{-v_1}) \geq n^{-v_2}) \leq cn^{-h_0}(1+t_0)^{m/2}, \quad n \in \mathbb{N}, \quad (\text{A.1})$$

where $c = c(v_1, v_2, m)$ does not depend on n or t_0 . Indeed, if $c_1 \geq 1$, the result follows directly from (A.1). If $c_1 \in (0, 1)$, let $\bar{v}_2 > v_2$ be such that it satisfies all hypotheses of the lemma. Namely, $\bar{v}_2 \in (0, 1)$, $v_1 \leq \bar{v}_2 + \frac{1}{2}$, $\bar{h}_0 := (\frac{m}{2} - 1)v_1 - m\bar{v}_2 > 0$. Then, by (A.1), for n such that $c_1 n^{-v_2} \geq n^{-\bar{v}_2}$,

$$\mathbb{P}(w_{t_0}(\hat{A}^n, n^{-v_1}) \geq c_1 n^{-v_2}) \leq c(1+t_0)^{\frac{m}{2}} n^{-\bar{h}_0}.$$

Moreover, $\bar{h}_0$ can be made arbitrarily close to h_0 by choosing $\bar{v}_2$ close to v_2 . This shows that the desired inequality holds for all large n , and by making $c = c(c_1, v_1, v_2, m)$ larger, for all $n \in \mathbb{N}$.

We now prove (A.1). As before, c denotes a positive constant whose value may change from line to line; here it may depend on (v_1, v_2, m) but not on n, t_0 . By (5.1),

$$\hat{A}^n(t) - \hat{A}^n(s) = n^{-1/2}(A^n(t) - A^n(s)) - n^{-1/2}\lambda^n(t-s).$$

Thus,

$$\mathbb{P}(w_{t_0}(\hat{A}^n, n^{-v_1}) \geq n^{-v_2}) \leq \mathbb{P}(A^n(t_0) > 2\lambda^n t_0) + \mathbb{P}(\Omega_+^n) + \mathbb{P}(\Omega_-^n), \quad (\text{A.2})$$

where

$$\begin{aligned} \Omega_+^n &= \left\{ A^n(t_0) \leq 2\lambda^n t_0, \exists s, t \in [0, t_0], 0 \leq t-s \leq n^{-v_1}, A^n(t) - A^n(s) \geq n^{\frac{1}{2}-v_2} + \lambda^n(t-s) \right\}, \\ \Omega_-^n &= \left\{ A^n(t_0) \leq 2\lambda^n t_0, \exists s, t \in [0, t_0], 0 \leq t-s \leq n^{-v_1}, A^n(t) - A^n(s) \leq -n^{\frac{1}{2}-v_2} + \lambda^n(t-s) \right\}. \end{aligned}$$

Consider the event Ω_+^n . Because $a^n(l)$ are the interarrival times of A^n , on this event, there must exist $l^0 \leq 2\lambda^n t_0$ and $R \leq n^{-v_1}$ such that

$$\sum_{l=l^0}^{l^0 + n^{\frac{1}{2}-v_2} + \lambda^n R} a^n(l) \leq R.$$

Recall that $\mathbb{E}\tilde{a}(l) = 1$. Letting $\bar{a}^n(l) = a^n(l) - (\lambda^n)^{-1}$, we have $\mathbb{E}\bar{a}^n(l) = 0$. Then, taking $r = \lambda^n R$, using $\lambda^n \leq c_1 n$ where $c_1 = \sup \frac{\lambda^n}{n} < \infty$, we have

$$\begin{aligned} \mathbb{P}(\Omega_+^n) &\leq \mathbb{P}\left(\exists j \leq 2\lambda^n t_0, \exists r \leq n^{-\nu_1} \lambda^n, \sum_{l=l^0+1}^{l^0+\frac{1}{n^{2-\nu_2}}+r} a^n(l) \leq (\lambda^n)^{-1} r\right) \\ &\leq \mathbb{P}\left(\exists j \leq 2c_1 n t_0, \exists r \leq c_1 n^{1-\nu_1}, \sum_{l=l^0+1}^{l^0+\frac{1}{n^{2-\nu_2}}+r} \bar{a}^n(l) \leq -(\lambda^n)^{-1} n^{\frac{1}{2}-\nu_2}\right). \end{aligned}$$

Let $M_{l^1}^n = \sum_{l=1}^{l^1} \bar{a}^n(l)$, $l^1 = 0, 1, 2, \dots$, and note that it is a martingale. For a real-valued function X on $\mathbb{Z}_+$ let

$$\text{osc}(X, l_1, l_2) = \max\{|X(l_3) - X(l_4)| : l_3, l_4 \in [l_1, l_2]\}, \quad 0 \leq l_1 \leq l_2.$$

Then, using $\frac{1}{2} - \nu_2 \leq 1 - \nu_1$ and denoting $\rho = \lfloor c_1 n^{1-\nu_1} \rfloor$,

$$\begin{aligned} \mathbb{P}(\Omega_+^n) &\leq \mathbb{P}(\exists j \leq 2c_1 n t_0, \text{osc}(M^n, l^0, l^0 + 2c_1 n^{1-\nu_1}) \geq c_1^{-1} n^{-\frac{1}{2}-\nu_2}) \\ &\leq \mathbb{P}(\exists j \in [0, 2c_1 n t_0] \cap \{0, \rho, 2\rho, \dots\}, \text{osc}(M^n, l^0, l^0 + \rho) \geq c_1^{-1} n^{-\frac{1}{2}-\nu_2} / 3) \\ &\leq 1 - (1 - \mathbb{P}(\|M^n\|_\rho \geq c_1^{-1} n^{-\frac{1}{2}-\nu_2} / 6))^{ct_0 n^{\nu_1}}. \end{aligned}$$

Because $n^{-1} \lambda^n \rightarrow \lambda > 0$, we have the lower bound $\lambda^n > cn$ for some $c > 0$ and all large n . Hence Burkholder's inequality shows that

$$\mathbb{E}(\|M^n\|_\rho)^m \leq c \mathbb{E}(\rho |\bar{a}^n(1)|^2)^{\frac{m}{2}} \leq cn^{(1-\nu_1)\frac{m}{2}} (\lambda^n)^{-m} \leq cn^{-(1+\nu_1)\frac{m}{2}}.$$

Hence, for any $a > 0$, $\mathbb{P}(\|M^n\|_\rho > a) \leq ca^{-m} n^{-(1+\nu_1)\frac{m}{2}}$, and, therefore,

$$\mathbb{P}(\Omega_+^n) \leq 1 - (1 - cn^{-(\nu_1-2\nu_2)\frac{m}{2}} t_0 n^{\nu_1})^{ct_0 n^{\nu_1}}.$$

Note that $\nu_1 > 2\nu_2$ by the assumption $h_0 > 0$. Now, if $\varepsilon \in (0, \frac{1}{2})$ and $a > 0$, then, with $c_2 = 2 \log 2$, $1 - (1 - \varepsilon)^a \leq 1 - e^{-c_2 a \varepsilon} \leq c_2 a \varepsilon$. This gives

$$\mathbb{P}(\Omega_+^n) \leq cn^{-(\nu_1-2\nu_2)\frac{m}{2}} t_0 n^{\nu_1} = ct_0 n^{-h_0}.$$

By a similar argument, the same estimate holds for $\mathbb{P}(\Omega_-^n)$. An application of (5.18) gives

$$\mathbb{P}(A^n(t_0) > 2\lambda^n t_0) \leq \mathbb{P}(\|\hat{A}^n\|_{t_0} > cn^{1/2}) \leq c \frac{\mathbb{E}(\|\hat{A}^n\|_{t_0}^m)}{n^{m/2}} \leq cn^{-\frac{m}{2}} (1 + t_0)^{\frac{m}{2}}.$$

Hence, by (A.2),

$$\mathbb{P}(w_{t_0}(\hat{A}^n, n^{-\nu_1}) \geq n^{-\nu_2}) \leq ct_0 n^{-h_0} + c(1 + t_0)^{\frac{m}{2}} n^{-\frac{m}{2}}.$$

It follows from $\nu_1, \nu_2 \in (0, 1)$ that $h_0 < \frac{m}{2}$. The result follows. $\square$

Proof of Lemma 5.9. Fix i, k . Denote $Y_t^n = \Lambda(S_{ik}^n, t)$. For any $u > 0$, if $Y_t^n \geq u$, then there must exist $0 \leq t_1 < t_2 \leq t$ with $t_2 - t_1 = u$ and $S_{ik}^n(t_2-) = S_{ik}^n(t_1)$; hence, by the Definition (5.1) of $\hat{S}_{ik}^n$,

$$\hat{S}_{ik}^n(t_1) - \hat{S}_{ik}^n(t_2-) = n^{-1/2} \mu_{ik}^n(t_2 - t_1) \geq c_2 n^{1/2} u,$$

for some constant c_2 that depends only on the sequence $\{\mu_{ik}^n\}$. Hence, $Y_t^n \leq 2c_2^{-1} n^{-1/2} \|\hat{S}_{ik}^n\|_t$. The first result now follows from (5.18) with $\kappa = 2$.

To prove the second statement, we will proceed in two steps. First, we will show that the number of interarrival times involved in the maximum is at most cn with probability going to one, then show that the maximum over cn variables has the right order of magnitude.

To simplify the notation, fix (i, k) and remove them from the notation of $S_{ik}^n, u_{ik}^n, \hat{\mu}_{ik}^n$, etc. The claim will be proved for $e_{\max}^n(t)$ defined as in (5.20), but without maximizing over (i, k) ; clearly, this is sufficient. Note that

$$S^n(t) = \sup \left\{ s \geq 0 : \sum_{l=1}^s u^n(l) \leq t \right\}.$$

Let

$$K^n = \inf \left\{ s \geq 0 : \sum_{l=1}^s u^n(l) \geq t_0 \right\}.$$

Then,

$$e_{\max}^n(t_0) \leq \sup \left\{ u^n(s) : \sum_{l=1}^{s-1} u^n(l) \leq t_0 \right\} \leq \sup \{ u^n(s) : s \leq K^n \}.$$

For any $c_2 \geq 0$,

$$\begin{aligned} \mathbb{P}(K^n \geq c_2 n) &\leq \mathbb{P} \left(\sum_{p=1}^{c_2 n} u^n(p) \leq t_0 \right) \\ &= \mathbb{P} \left(\sum_{p=1}^{c_2 n} \check{u}(p) \leq \mu^n t_0 \right) \\ &= \mathbb{P} \left(\frac{1}{c_2 n} \sum_{p=1}^{c_2 n} \check{u}(p) \leq \frac{\mu t_0}{c_2} + \frac{\hat{\mu}^n t_0}{c_2} n^{-1/2} \right). \end{aligned}$$

If $c_2 > \mu t_0$, by the law of large numbers, the RHS converges to 0 as $n \rightarrow \infty$. Next, by independence of service times, for any $c > 0$,

$$U_n := \mathbb{P} \left(\max_{l \leq c_2 n} u^n(l) \geq c n^{\bar{a}-1} \right) = 1 - \mathbb{P} \left(u^n(1) \leq c n^{\bar{a}-1} \right)^{c_2 n}.$$

Using $1 - (1-x)^n \leq c_3 n x$, letting $c_4 = c_2 c_3$ and denoting $\bar{\varepsilon} = \mathbf{m} - 2 > 0$, we obtain

$$\begin{aligned} U_n &\leq c_4 n \mathbb{P}(u^n(1) \geq c n^{\bar{a}-1}) \\ &\leq c_4 n \mathbb{P}(\check{u}(1) \geq c \mu^n n^{\bar{a}-1}) \\ &= c_4 n \mathbb{P}((\check{u}(1))^{2+\bar{\varepsilon}} \geq (c \mu^n n^{\bar{a}-1})^{2+\bar{\varepsilon}}) \\ &\leq c_4 n \frac{\mathbb{E}[(\check{u}(1))^{2+\bar{\varepsilon}}]}{(c \mu^n n^{\bar{a}-1})^{2+\bar{\varepsilon}}}, \end{aligned}$$

where the last inequality uses Assumption 2.8. Next, by the definition in (2.31) of $\bar{a}$, we have that $\bar{a} > \frac{1}{2} - \frac{\bar{\varepsilon}}{4(\bar{\varepsilon}+2)}$, hence

$$\begin{aligned} (2 + \bar{\varepsilon})(1 + \bar{a} - 1) &= \bar{a}(2 + \bar{\varepsilon}) \\ &\geq (2 + \bar{\varepsilon}) \left(\frac{1}{4} + \frac{1}{4 + 2\bar{\varepsilon}} \right) \\ &= \frac{1}{2} + \frac{\bar{\varepsilon}}{4} + \frac{2 + \bar{\varepsilon}}{4 + 2\bar{\varepsilon}} = 1 + \frac{\bar{\varepsilon}}{4}. \end{aligned}$$

Thus,

$$n \frac{\mathbb{E}[(\check{u}(1))^{2+\bar{\varepsilon}}]}{(c \mu^n n^{\bar{a}-1})^{2+\bar{\varepsilon}}} = O(n^{-\frac{\bar{\varepsilon}}{4}}).$$

Hence, for any $c > 0$, there exists c_2 such that

$$\begin{aligned} \mathbb{P}(e_{\max}^n \geq c n^{\bar{a}-1}) &\leq \mathbb{P}(K^n \geq c_2 n) + \mathbb{P}(e_{\max}^n \geq c n^{\bar{a}-1}, K^n \leq c_2 n) \\ &\leq \mathbb{P}(K^n \geq c_2 n) + \mathbb{P} \left(\max_{l \leq c_2 n} u^n(l) \geq c n^{\bar{a}-1} \right), \end{aligned}$$

and both terms have been shown to converge to zero. $\square$

Proof of Proposition 5.3 (Continued from Section 5.2.2). Part 1 of the proof of this proposition appears in Section 5.2.2; here, we provide the remaining parts.

Part 2. Consider the case where $\mathbf{T}_1$ is applicable when the workload is sufficiently large. This covers case 2(d) of Theorem 2.13, in which the policy $\mathbf{T}_2\mathbf{T}_1$ applies (none of our proposed policies implement $\mathbf{T}_1$ as a single-mode policy). The proof given in part 1 is applicable, for the following reasons. In the first step, during the analyzed time interval $[\tau, t]$, one has $\hat{X}_{2,2}^n \geq K > \hat{\Theta}^n$, and, therefore, there is no difference between how $\mathbf{P}$ and $\mathbf{T}_1$ behave during this interval. In addition, the workload is sampled at each arrival/service, which means that

$$T_{22}^n(t) - T_{22}^n(\tau) \geq t - \tau - e_{\max}^n.$$

In the second step, the argument given for nonidling of both servers during $[\sigma, t]$ again holds here, similarly to part 1, upon noticing that, because σ corresponds to an arrival, it is a sampling time, and so the current mode is either already the high mode or switches to the high mode at that time. (It is possible that, if σ is a mode switching time, then server 1 was idle just before σ . But, if so, it starts to serve class 1 at σ .) The remaining details need no adaptation.

Part 3. Finally, consider the policies that employ T_2 for high workload levels, namely, the policies T_2 , T_2T_2 , and T_1T_2 , covering all remaining cases of Theorem 2.13. In these cases, the stronger moment assumption is in force, and the goal is to prove that there exists $\varepsilon_0 > 0$ such that $\mathbb{E}[(\hat{W}_t^n)^{1+\varepsilon_0}] \leq \text{pol}(t)$ for all n and t , for some polynomial pol . As before, let ξ^A or ξ^H be in canonical form, in the single-mode and, respectively, dual-mode case. Fix a constant $K > z^*$ in the dual-mode case and $K = 1$ in the single-mode case. Given t , consider the event $\hat{W}_t^n > K$. Let

$$\tau_1 = \tau_1^n = \sup\{s \in [0, t] : \hat{W}_s^n \leq K\}.$$

Then, $\hat{W}_{\tau_1}^n \leq K + 1$, and, moreover, $\hat{W}_t^n \geq K$ during $[\tau_1, t]$. Although this lower bound on the workload is sufficiently large to guarantee that $\hat{W}_t^n$ corresponds to the higher workload mode, it is possible that at time τ_1 , the current mode variable still equals the lower workload mode. We argue that the time it takes to switch to the upper workload mode is bounded by $2e_{\max}^n$ in both T_1T_2 and T_2T_2 . In the former case, the current mode switches as soon as there is a new arrival or departure. In the latter case, one possibly has to complete the service of a job at server 2, which is server $k_1(\xi^L)$. It is possible that there are no jobs allowed to be routed to server 2 at time τ_1 . Wait for an arrival of class 1 job, and service completion of this job at server 2. At this time, it is guaranteed that mode has switched to ξ^H if it was not ξ^H earlier. Thus, if we let

$$\tau_2 = \inf\{s \geq \tau_1 : \text{MODE}^n(s) = \xi^H\} \wedge t,$$

we have $\tau_2 - \tau_1 \leq 2e_{\max}^n \wedge t$.

According to the rules of T_2 , server 2 must be busy throughout the interval $[\tau_2, t]$. Thus, $\hat{I}_2^n(t) = \hat{I}_2^n(\tau_2)$. Hence, by (5.6), and using $\hat{I}_k^n[a, b] \leq n^{1/2}(b - a)$ (by (5.2)), we have

$$\begin{aligned} \hat{W}_t^n &= \hat{W}_{\tau_1}^n + \hat{F}^n[\tau_1, t] + \hat{L}^n[\tau_1, t] \\ &= \hat{W}_{\tau_1}^n + \hat{F}^n[\tau_1, t] + \beta_1 \hat{I}_1^n[\tau_1, t] + \beta_2 \hat{I}_2^n[\tau_1, t] \\ &\leq K + 1 + 2\|\hat{F}^n\|_t + cn^{1/2}e_{\max}^n + \beta_1 \hat{I}_1^n[\tau_2, t], \end{aligned}$$

holds on the event $\hat{W}_t^n > K$. On the complementary event, $\hat{W}_t^n \leq K$. Use (5.5) and (5.18) to obtain the bound $\mathbb{E}[\|\hat{F}^n\|_t^{1+\varepsilon_0}] \leq \{\mathbb{E}[\|\hat{F}^n\|_t^2]\}^{(1+\varepsilon_0)/2} \leq c(1+t)^{(1+\varepsilon_0)/2}$. Use Lemma 5.9 to bound the second moment of $n^{1/2}e_{\max}^n$ by $c(1+t)$. By Min-kowski's inequality and the inequality $(a+b)^{1+\varepsilon_0} \leq 4a^{1+\varepsilon_0} + 4b^{1+\varepsilon_0}$, which holds for $a, b \geq 0$, $\varepsilon_0 \in (0, 1)$, this yields

$$\mathbb{E}[(\hat{W}_t^n)^{1+\varepsilon_0}] \leq c(1+t)^{(1+\varepsilon_0)/2} + c\mathbb{E}[\mathbb{1}_{\{\hat{W}_t^n > K\}} \hat{I}_1^n[\tau_2, t]^{1+\varepsilon_0}].$$

Hence, it suffices to bound the last term above by a polynomial in t .

To this end, let $\tau_3 = \inf\{s \geq \tau_2 : \hat{X}_1^n(s) \geq \hat{\Theta}^n\} \wedge t$. Then,

$$\begin{aligned} \mathbb{E}[\mathbb{1}_{\{\hat{W}_t^n > K\}} \hat{I}_1^n[\tau_2, t]^{1+\varepsilon_0}] &\leq 4\Delta_1^n + 4\Delta_2^n, \quad \text{where} \\ \Delta_1^n &= \mathbb{E}[\mathbb{1}_{\{\hat{W}_t^n > K\}} \hat{I}_1^n[\tau_2, \tau_3]^{1+\varepsilon_0}] \\ \Delta_2^n &= \mathbb{E}[\mathbb{1}_{\{\hat{W}_t^n > K, \tau_3 < t\}} \hat{I}_1^n[\tau_3, t]^{1+\varepsilon_0}]. \end{aligned}$$

To bound Δ_1^n , consider the event $\{\hat{W}_t^n > K, \tau_3 > \tau_2\}$. On it, during the interval $[\tau_2, \tau_3]$, one has $\hat{X}_1^n < \hat{\Theta}^n$; hence, by the rules of T_2 , server 2 prioritizes class 2, except possibly it completes a service that started when $\hat{X}_1^n \geq \hat{\Theta}^n$. Moreover, $\hat{W}_t^n \geq K$ during the same interval, by which we know that there are multiple class-2 jobs in the system, and, thus, server 2 gives no service to class 1, with the only exception of service to a job that started when $\hat{X}_1^n \geq \hat{\Theta}^n$. Hence, the departure process associated with activity (1, 2) increases by at most one during this interval that is, $0 \leq D_{12}^n[\tau_2, \tau_3] \leq 1$. Recalling the definition of $\hat{S}^n$ in (5.1), this can be expressed as $0 \leq e_1^n[\tau_2, s] \leq n^{-1/2}$, $s \in [\tau_2, \tau_3]$, where

$$e_1^n(s) := n^{-1/2}D_{12}^n(s) = \hat{S}_{12}^n(T_{12}^n(s)) + n^{-1/2}\mu_{12}^n T_{12}^n(s).$$

Using this in (5.4) gives, for $s \in [\tau_2, \tau_3]$,

$$\hat{X}_1^n(s) = \hat{X}_1^n(\tau_2) + \hat{f}_1^n[\tau_2, s] - e_1^n[\tau_2, s] + n^{1/2}\lambda_1(s - \tau_2) - n^{1/2}\mu_{11}^n T_{11}^n[\tau_2, s],$$

where

$$\hat{f}_1^n(s) = \hat{A}_1^n(s) - \hat{S}_{11}^n(T_{11}^n(s)) + (\hat{\lambda}_1^n - \hat{\mu}_{11}^n) T_{11}^n(s).$$

Next, by (2.4),

$$T_{11}^n[\tau_2, s] = (s - \tau_2) - I_1^n[\tau_2, s] - T_{21}^n[\tau_2, s].$$

However, activity (2, 1) is not in use by T_2 . Denoting $c_1 = \lambda_1 - \mu_{11}$ and $g(s) = c_1 s$, this yields

$$\hat{X}_1^n(s) = \hat{X}_1^n(\tau_2) + \hat{f}_1^n[\tau_2, s] - e_1^n[\tau_2, s] + n^{1/2}g[\tau_2, s] + \mu_{11}\hat{I}_1^n[\tau_2, s], \quad s \in [\tau_2, \tau_3]. \quad (\text{A.3})$$

Because T_2 is used after the update of modes, it must be true that $c_1 = \lambda_1 - \mu_{11} > 0$.

We use the following property of the Skorohod map. Let $\eta = \Gamma_2[\psi]$. Then, by (5.7), $\eta(s) = \sup_{0 \leq \theta \leq s} [\psi(\theta)^-]$. Assume that $\psi = \psi_1 + \psi_2$ where $\psi_2 \geq 0$. Then,

$$\begin{aligned} \eta(s) &= \Gamma_2[\psi_1 + \psi_2](s) = \sup_{\theta \in [0, s]} [(\psi_1(\theta) + \psi_2(\theta))^-] \\ &\leq \sup_{\theta \in [0, s]} [\psi_1(\theta)^-] \\ &\leq \|\psi_1\|_s. \end{aligned}$$

This property is used as follows. On the interval $[\tau_2, \tau_3]$, $\hat{I}_1^n$ can increase only at times when $\hat{X}_1^n = 0$; and the latter process is nonnegative. This shows that $\mu_{11}\hat{I}_1^n$ serves as the Skorohod term in (A.3) on that interval, and, thus, by the non-negativity of $\hat{X}_1^n(\tau_2)$ and c_1 , and the bound $|e_1^n[\tau_2, s]| \leq n^{-1/2}$, this shows that

$$\mu_{11}\hat{I}_1^n[\tau_2, \tau_3] \leq \sup_{\theta \in [\tau_2, \tau_3]} |\hat{f}_1^n[\tau_2, \theta]| + \sup_{\theta \in [\tau_2, \tau_3]} |e_1^n[\tau_2, \theta]| \leq 2\|\hat{f}_1^n\|_t + n^{-1/2}.$$

Hence, $\Delta_1^n \leq c\{\mathbb{E}[\|\hat{f}_1^n\|_t^2]\}^{(1+\varepsilon_0)/2} + 1 \leq c(1+t)$.

Next, we bound Δ_2^n . To this end, consider the event $\Omega^n := \{\hat{W}_t^n > K, \tau_3 < t, \hat{I}_1^n(t) > \hat{I}_1^n(\tau_3)\}$. Because by its definition, $\hat{I}_1^n(t) \leq n^{1/2}t$, we have

$$\Delta_2^n \leq \mathbb{P}(\Omega^n)(n^{1/2}t)^{1+\varepsilon_0}.$$

On the event Ω^n , let

$$\tau_5 = \inf\{s > \tau_3 : \hat{X}_1^n(s) = 0\}, \quad \tau_4 = \sup\{s < \tau_5 : \hat{X}_1^n(s) \geq \hat{\Theta}^n\}.$$

Then, on Ω^n , it must hold that $\tau_3 \leq \tau_4 < \tau_5 \leq t$. The arguments that led to (A.3) are valid for the time interval $[\tau_4, \tau_5]$. As a result,

$$\hat{X}_1^n[\tau_4, \tau_5] = \hat{f}_1^n[\tau_4, \tau_5] - e_1^n[\tau_4, \tau_5] + n^{1/2}g[\tau_4, \tau_5] + \mu_{11}\hat{I}_1^n[\tau_4, \tau_5],$$

with $|e_1^n[\tau_4, s]| \leq n^{-1/2}$ for all $s \in [\tau_4, \tau_5]$. Now, $\hat{I}_1^n$ remains flat on the interval $[\tau_4, \tau_5]$, and, moreover, $\hat{X}_1^n(\tau_4) = \hat{\Theta}^n - n^{-1/2}$ and $\hat{X}_1^n(\tau_5) = 0$. This gives

$$\hat{f}_1^n[\tau_4, \tau_5] + n^{1/2}g[\tau_4, \tau_5] \leq -\hat{\Theta}^n + 2n^{-1/2}.$$

Given any $\delta > 0$, using the nonnegativity of the second term on the LHS, in the case that $\tau_5 - \tau_4 \leq n^{-\delta}$, one must have $\hat{f}_1^n[\tau_4, \tau_5] \leq -\hat{\Theta}^n/2$. On the other hand, in the case $\tau_5 - \tau_4 > n^{-\delta}$, one must have $\hat{f}_1^n[\tau_4, \tau_5] + c_1 n^{1/2}n^{-\delta} < 0$. As a result,

$$\mathbb{P}(\Omega^n) \leq p_1^n + p_2^n := \mathbb{P}(w_t(\hat{f}_1^n, n^{-\delta}) \geq \hat{\Theta}^n/2) + \mathbb{P}(2\|\hat{f}_1^n\|_t > c_1 n^{\frac{1}{2}-\delta}).$$

We have by (2.32) $\hat{\Theta}^n = n^{-\frac{1}{2}}\lceil n^{\bar{a}} \rceil$, where we recall that $\bar{a} < \frac{1}{2}$. Thus, by Lemma 5.7, with $\nu_1 = \delta$, $\nu_2 = \frac{1}{2} - \bar{a}$, one has, for any $\varepsilon_1 > 0$,

$$p_1^n \leq c(1+t)^{\frac{\mathbf{m}}{2}} n^{-h_0+\varepsilon_1},$$

where

$$h_0 = \left(\frac{\mathbf{m}}{2} - 1\right)\delta - \mathbf{m}\left(\frac{1}{2} - \bar{a}\right), \quad \delta \leq 1 - \bar{a}.$$

Next, by (5.18) and Chebychev's inequality,

$$p_2^n \leq c(1+t)^{\frac{\mathbf{m}}{2}} n^{-\frac{\mathbf{m}}{2}+\delta\mathbf{m}}.$$

As a result, we have

$$\Delta_2^n \leq c(n^{\zeta_1(\delta, \varepsilon_0, \varepsilon_1)} + n^{\zeta_2(\delta, \varepsilon_0)})(1+t)^{\frac{\mathbf{m}}{2}+1+\varepsilon_0},$$

where

$$\zeta_1(\delta, \varepsilon_0, \varepsilon_1) = -\left(\frac{\mathbf{m}}{2} - 1\right)\delta + \mathbf{m}\left(\frac{1}{2} - \bar{a}\right) + \frac{1}{2} + \frac{\varepsilon_0}{2} + \varepsilon_1, \quad \zeta_2(\delta, \varepsilon_0) = -\frac{\mathbf{m}}{2} + \delta\mathbf{m} + \frac{1}{2} + \frac{\varepsilon_0}{2}.$$

By our assumptions, we have $\mathbf{m} > \mathbf{m}_0 = \frac{1}{2}(5 + \sqrt{17})$ and $\bar{a} > \frac{1}{2} - \frac{\mathbf{m}^2 - 5\mathbf{m} + 2}{2\mathbf{m}(3\mathbf{m} - 2)}$. Using this, a calculation shows that with the choice $\delta = \mathbf{m}/(3\mathbf{m} - 2) \in (1/3, 1/2)$, one has $\zeta_1(\delta, 0, 0) \vee \zeta_2(\delta, 0) < 0$. It follows that there exist $\varepsilon_0 > 0$ and $\varepsilon_1 > 0$ for which $\zeta_1(\delta, \varepsilon_0, \varepsilon_1) \vee \zeta_2(\delta, \varepsilon_0) < 0$. Therefore, $\Delta_2^n \leq c(1 + t)^{\frac{\mathbf{m}}{2} + 1 + \varepsilon_0}$ and the proof is complete. $\square$

Proof of Lemma 5.10 (Continued from Section 5.2.3). The proof of the lemma in case 1(a) appears in Section 5.2.3; here, we cover the remaining cases.

• **Case 1(b):** In this case, we use the T_2 policy, which is single mode. Between σ and τ , $\hat{X}_p^n(t) \geq \hat{\Theta}^n$. Thus, both servers give priority to the HPC at all time in $[\sigma, \tau]$. It is possible that the dual-activity server is occupied with the low-priority class at σ , but as soon as the current job is served, the high-priority class gets priority on one server and dedication by the other server. By definition of $e_{\max}^n$, we have for any $k \in \{1, 2\}$,

$$\int_{\sigma}^{\tau} \Xi_{pk}^n(t) dt \geq \tau - \sigma - e_{\max}^n \mathbb{1}_{k=k_2(\xi^A)}.$$

Thus,

$$\int_{\sigma}^{\tau} \left(\lambda_p - \sum_k \mu_{pk} \Xi_{pk}^n(t) \right) dt \leq (\tau - \sigma) \left(\lambda_p - \sum_k \mu_{pk} \right) + \mu_{pk_2(\xi^A)} e_{\max}^n.$$

Finally, $\lambda_p - \sum_k \mu_{pk} < 0$ by (3.1) and $\min_i \lambda_i > 0$.

• **Case 2(a):** In this case, we use the PP policy, which only changes mode at the completion of a service at the single-activity server. This is precisely the server that gives priority to the high-priority class after switching mode because we are in an SS case. Because there are always HPC jobs to serve between σ and τ , we claim that excluding an initial period, at any given time the HPC is being served by at least one server, as discussed in Section 5.1.3. We now discuss how long this initial period can be.

The only way some service is lost is if there were no high-priority jobs at some point in the system, both servers become busy with low-priority jobs, and σ occurs before the service completion of those jobs. After completion of those two jobs, the HPC keeps priority on at least one server, regardless of switching of the current mode. For the initial jobs, with respect to $\text{MODE}^n(\sigma)$, either the service ends first at the dual-activity server or the service ends first at the single-activity server.

In the first case, a service of the HPC job starts at the dual-activity server because it has priority there until either τ is reached or a mode switch occurs. In the second case, the available server is currently dedicated to the LPC. Because this corresponds to a service completion at the single-activity server, the workload is sampled, and there could also be a switching of modes. If the mode changes, then this server becomes dual activity, and the HPC has priority there. If there is no mode switch, the service of another LPC job starts. LPC jobs are served at this server until the mode switches.

Thus, if the service ends at the dual-activity server before the mode switches, then an HPC job will start there. If the mode switches before the service ends at the current dual-activity server, then the current single-activity server becomes dual activity and an HPC job begins service there.

In either case, the HPC is served by at least one server after that time because it has priority on the dual-activity server, and the dual-activity server is always available when switching modes. In addition, the time it takes for the HPC to begin service is smaller than the service of the job present at the dual-activity server at time σ , which is smaller than $e_{\max}^n$. Hence,

$$\int_{\sigma}^{\tau} (\mu_{p1} \Xi_{p1}^n(t) + \mu_{p2} \Xi_{p2}^n(t)) dt \geq \min_k \mu_{pk} (\tau - \sigma) - \max_k \mu_{pk} e_{\max}^n.$$

This yields

$$\int_{\sigma}^{\tau} \left(\lambda_p - \sum_k \mu_{pk} \Xi_{pk}^n(t) \right) dt \leq (\tau - \sigma) \left(\lambda_p - \min_k \mu_{pk} \right) + \max_k \mu_{pk} e_{\max}^n.$$

In addition, $\lambda_p - \min_k \mu_{pk} < 0$. Putting ξ^L in canonical form forces in this case $\frac{\lambda_1}{\alpha_1} > \beta_1$ and $p = 2$ because $i_1(\xi^L) = p$. In addition, by (3.6), either $\frac{\lambda_1}{\alpha_1} > \beta_1 \vee \beta_2$ or $\frac{\lambda_2}{\alpha_2} > \beta_1 \vee \beta_2$. Thus, $\frac{\lambda_1}{\alpha_1} > \max_k \beta_k$, which implies $\frac{\lambda_2}{\alpha_2} < \min_k \beta_k$ by (3.1).

• **Case 2(b):** In this case, we use the T_2T_2 policy. Because the number of HPC jobs stays above Θ^n throughout the period $[\sigma, \tau]$, both servers only take new jobs from the HPC, regardless of the mode. Similarly to 1(b), there is at most one job served at either activity before the HPC gets served. Hence, for any $k \in \{1, 2\}$,

$$\int_{\sigma}^{\tau} \Xi_{pk}^n(t) dt \geq \tau - \sigma - e_{\max}^n.$$

Thus,

$$\int_{\sigma}^{\tau} \left(\lambda_p - \sum_k \mu_{pk} \Xi_{pk}^n(t) \right) dt \leq (\tau - \sigma) \left(\lambda_p - \sum_k \mu_{pk} \right) + \sum_k \mu_{pk} e_{\max}^n.$$

We obtain the result similarly as before because $\lambda_p - \sum_k \mu_{pk} < 0$.

• **Cases 2(c) and 2(d):** The HPC is single activity in one mode and dual activity in the other. Recall that this case is CS. This means the dual activity server stays the same, regardless of switching modes. When a T_1 rule is applied, the dual-activity server prioritizes the single-activity class as long as there are more than Θ^n jobs of this class. The HPC is single activity when we apply the T_1 rule. Similarly, under a T_2 rule, the dual-activity server prioritizes the dual-activity class as long as there are more than Θ^n jobs of this class and the HPC is dual activity when we apply the T_2 rule. Once again, put ξ^L in canonical form, server 2 is the dual-activity server in both modes. Regardless of which rule is used, as long as the number of HPC jobs is above Θ^n , server 2 only takes new jobs from the HPC. As before, we have to exclude a residual service of a low-priority job that started before σ . Hence,

$$\int_{\sigma}^{\tau} \Xi_{p2}^n(t) dt \geq \tau - \sigma - e_{\max}^n.$$

It remains to show that the activity processing HPC jobs throughout the period is enough to deplete them. Because of the canonical form of ξ^L , $\frac{\lambda_1}{\alpha_1} > \beta_1$ and, thus, by (3.5), $\frac{\lambda_1}{\alpha_1} < \beta_2$. Moreover, $\frac{\lambda_2}{\alpha_2} < \beta_2$ by (3.1). Whichever the HPC is, server 2 has enough capacity to deplete it. In other words, $\lambda_p - \mu_{p2} < 0$ and

$$\int_{\sigma}^{\tau} \left(\lambda_p - \sum_k \mu_{pk} \Xi_{pk}^n(t) \right) dt \leq (\tau - \sigma)(\lambda_p - \mu_{p2}) + \mu_{p2} e_{\max}^n. \quad \square$$

Proof of Lemma 5.12. The proof follows reasoning similar to the proof of Proposition 5.4. Fix δ . Let us introduce

$$\sigma_r^n = \sup \left\{ t \leq \tau_r^n : \hat{X}_p^n(t) \geq \hat{\Theta}^n, \text{ or } \hat{X}_q^n(t) \leq \frac{\hat{\Theta}^n}{2} \right\}.$$

In cases 2(b), (c), and (d), introduce

$$\begin{aligned} \tilde{\tau}^n &= \inf \left\{ t \geq 0 : \hat{X}_q^n(t) \geq \frac{z^* \alpha_q}{2} \right\}, \\ \tilde{\sigma}^n &= \sup \left\{ t \leq \tilde{\tau}^n : \hat{X}_q^n(t) \leq \frac{z^* \alpha_q}{4} \right\}. \end{aligned}$$

To simplify, we write throughout this proof

$$\rho^n = \rho, \quad \tau_r^n = \tau, \quad \tilde{\tau}^n = \tilde{\tau}, \quad \tilde{\sigma}^n = \tilde{\sigma} \text{ and } \sigma_r^n = \sigma.$$

We briefly describe the interaction between the times we just introduced. Under the state-space collapse, $\tilde{\tau}$ has to occur before the first time the current mode switches. The times σ and $\tilde{\sigma}$ allow us to have some knowledge about the state of queue lengths, uniformly over an interval. The first step of the proof of (5.26) is establishing that

$$\mathbb{P}(\tau_c^n \geq t_0, \tilde{\tau} \leq \rho) \rightarrow 0, \quad (\text{A.4})$$

holds in cases 2(b), (c), and (d). In those cases, under the event $\{\tau_c^n \geq t_0\} \cap \{\tilde{\tau} \leq \rho\}$, for any $t \leq \rho$,

$$\text{MODE}^n(t) = \text{MODE}^n(0) = \xi^L. \quad (\text{A.5})$$

In addition, in case 1(b), for any $t \leq \rho$

$$\text{MODE}^n(t) = \xi^A, \quad (\text{A.6})$$

which means (A.4) is not needed in this case.

We now proceed with the proof of (A.4). On $\{\tau_c^n \geq t_0, \tilde{\tau} \leq \rho\}$, the following hold:

(i) The number of HPC job is below $\hat{\Theta}^n$, and there are LPC jobs in the system during $[\tilde{\sigma}, \tilde{\tau}]$:

$$\sup_{t \in [\tilde{\sigma}, \tilde{\tau}]} \hat{X}_p^n(t) < \hat{\Theta}^n \text{ and } \inf_{t \in [\tilde{\sigma}, \tilde{\tau}]} \hat{X}_q^n(t) > 0.$$

(ii) The current mode is the same as in the initial state: for any $t \in [\tilde{\sigma}, \tilde{\tau}]$,

$$\text{MODE}^n(t) = \text{MODE}^n(0) = \xi^L.$$

(iii) In addition, the number of LPC jobs must have grown during that time:

$$\hat{X}_q^n(\tilde{\tau}) - \hat{X}_q^n(\tilde{\sigma}) \geq \frac{z^* \alpha_q}{4}.$$

(i) comes from the definition of ρ , $\tilde{\sigma}$ and $\tilde{\tau}$. (ii) comes from the fact that the number of HPC jobs is bounded by $2\Theta^n$ under $\{\tau_c^n \geq t_0\}$ and the number of LPC jobs is bounded by $\frac{z^* \alpha_q}{2}$ before $\tilde{\tau}$ and the workload cannot cross above z^* . (iii) comes from the definition of $\tilde{\tau}$ and $\tilde{\sigma}$. Because of (ii), only the rule in the lower-workload mode is in use. When using a T_1 rule, both

servers take new jobs from the LPC because of (i), and $\lambda_q - \sum_k \mu_{qk} < 0$. When using a T_2 rule, again because of (i), the dual-activity server takes new jobs from the LPC. Because ξ^L is in canonical form $\frac{\lambda_1}{\alpha_1} > \beta_1$, $k_2(\xi^L) = 2$ and $q = 2$ because of the T_2 rule. This means that $\frac{\lambda_2}{\alpha_2} < \beta_2$ and, thus, $\lambda_q - \mu_{qk_2(\xi^L)} < 0$. Hence, with (5.19), regardless of the case there exists $c > 0$ such that

$$\hat{X}_q^n(\tilde{\tau}) \leq \hat{X}_q^n(\tilde{\sigma}) + \hat{f}_q^n[\tilde{\sigma}, \tilde{\tau}] - c\sqrt{n}(\tilde{\tau} - \tilde{\sigma}) + \sum_k \mu_{qk} e_{\max}^n.$$

From this, we obtain two bounds:

$$\begin{aligned} \hat{X}_q^n(\tilde{\tau}) - \hat{X}_q^n(\tilde{\sigma}) &\leq w_{t_0}(\hat{f}_q^n, \tilde{\tau} - \tilde{\sigma}) - c\sqrt{n}(\tilde{\tau} - \tilde{\sigma}) + \sum_k \mu_{qk} e_{\max}^n \\ \hat{X}_q^n(\tilde{\tau}) - \hat{X}_q^n(\tilde{\sigma}) &\leq 2\|\hat{f}_q^n\|_{t_0} - c\sqrt{n}(\tilde{\tau} - \tilde{\sigma}) + \sum_k \mu_{qk} e_{\max}^n. \end{aligned}$$

We now prove (A.4): let $r_n = \frac{\log(n+1)}{\sqrt{n}}$. Because of (iii),

$$\begin{aligned} \mathbb{P}(\tau_c^n \geq t_0, \tilde{\tau} \leq \rho) &\leq \mathbb{P}\left(\tilde{\tau} \leq \rho, \hat{f}_q^n[\tilde{\sigma}, \tilde{\tau}] - c\sqrt{n}(\tilde{\tau} - \tilde{\sigma}) + \sum_k \mu_{qk} e_{\max}^n \geq \frac{z^* \alpha_q}{4}\right) \\ &\leq \mathbb{P}\left(w_{t_0}(\hat{f}_q^n, \tilde{\tau} - \tilde{\sigma}) - c\sqrt{n}(\tilde{\tau} - \tilde{\sigma}) + \sum_k \mu_{qk} e_{\max}^n \geq \frac{z^* \alpha_q}{4}, \tilde{\tau} - \tilde{\sigma} \leq r^n\right) \\ &\quad + \mathbb{P}\left(2\|\hat{f}_q^n\|_{t_0} - c\sqrt{n}(\tilde{\tau} - \tilde{\sigma}) + \sum_k \mu_{qk} e_{\max}^n \geq \frac{z^* \alpha_q}{4}, \tilde{\tau} - \tilde{\sigma} > r^n\right) \\ &\leq \mathbb{P}\left(w_{t_0}(\hat{f}_q^n, r_n) + \sum_k \mu_{qk} e_{\max}^n \geq \frac{z^* \alpha_q}{4}\right) + \mathbb{P}\left(2\|\hat{f}_q^n\|_{t_0} + \sum_k \mu_{qk} e_{\max}^n \geq c \log(n+1)\right). \end{aligned}$$

By Lemma 5.9, $e_{\max}^n$ is smaller than $n^{\bar{a}-1} = o(1)$. By Remark 5.8, $\hat{f}_q^n$ are C-tight and $\|\hat{f}_q^n\|_{t_0}$ are tight RVs. Both terms must then converge to zero. This concludes the proof of (A.4) in all relevant cases for this lemma.

Case 2(c), T_1 rule:

By splitting the integral that defines $\bar{R}^n$, and using (5.24) and (A.5),

$$\begin{aligned} \bar{R}_\rho^n &\leq \int_0^\rho \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n, \tau_c^n \geq t_0, \tilde{\tau} \geq \rho} d\hat{L}_t^n + \int_0^\rho \mathbb{1}_{\tau_c^n < t_0} + \mathbb{1}_{\tilde{\tau} \leq \rho, \tau_c^n \geq t_0} d\hat{L}_t^n \\ &\leq \int_0^\rho \mathbb{1}_{\hat{X}_q^n(t) \geq \hat{\Theta}^n, \text{MODE}^n(t) = \xi^L} d\hat{L}_t^n + \int_0^\rho \mathbb{1}_{\tau_c^n < t_0} + \mathbb{1}_{\tilde{\tau} \leq \rho, \tau_c^n \geq t_0} d\hat{L}_t^n. \end{aligned}$$

The first term is zero by definition of the T_1 rule because q is the dual-activity class. The second term goes to zero by Proposition 5.4 and (A.4). This concludes the proof of (5.26) in case 2(c).

Case 1(b), 2(b), and 2(d), T_2 rule:

We now deal with all cases that use a T_2 rule in lower workload at the same time thanks to (A.5)/(A.6). Recall that $p = 1 = i_2(\xi^L)$ and $k_1(\xi^L) = 1$. By splitting the integral, we obtain

$$\bar{R}_\rho^n = \int_0^\rho \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n, \tau_c^n \geq t_0} d\hat{L}_t^n + \int_0^\rho \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n, \tau_c^n < t_0} d\hat{L}_t^n.$$

By (5.2), $\hat{L}^n = \sum_k \beta_k \hat{I}_k^n$. By definition of the T_2 rule, as long as there are at least two customers in the system, the dual-activity server cannot idle. With ξ^L in canonical form, server 2 is the dual-activity server, so under the event $\{\tau_c^n \geq t_0\} \cap \{\tilde{\tau} \geq \rho\}$,

$$\int_0^\rho \mathbb{1}_{\hat{W}_t^n \geq 2(\alpha_1 \wedge \alpha_2)^{-1} n^{-1/2}, \tau_c^n \geq t_0, \tilde{\tau} \geq \rho} d\hat{I}_2^n(t) = 0.$$

Because $\int_0^\rho \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n, \tau_c^n < t_0} d\hat{L}_t^n \Rightarrow 0$ by Proposition 5.4, we are left to deal with

$$\int_0^\rho \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n, \tau_c^n \geq t_0, \tilde{\tau} \geq \rho} d\hat{L}_t^n = \beta_1 \int_0^\rho \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n, \tau_c^n \geq t_0, \tilde{\tau} \geq \rho} d\hat{I}_1^n(t).$$

We introduce the following times:

$$\bar{\tau}^n = \inf \left\{ t \geq 0 : \int_0^t \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2} \right\},$$

and

$$\bar{\sigma}^n = \sup \{ t \leq \bar{\tau}^n : \hat{X}_2^n(t) = 0 \}.$$

We want to show that

$$\mathbb{P}(\rho \geq \bar{\tau}) \rightarrow 0.$$

On the event $\{\rho \geq \bar{\tau}\}$, we have:

- First, by definition of $\bar{\sigma}$,

$$\inf_{t \in (\bar{\sigma}, \bar{\tau})} \hat{X}_2^n(t) > 0.$$

- Second, by definition of ρ ,

$$\sup_{t \in [0, \bar{\tau}]} \hat{X}_1^n(t) < \hat{\Theta}^n.$$

- Finally, by definition of $\bar{\sigma}$,

$$\hat{X}_2^n(\bar{\sigma}) \leq n^{-1/2}.$$

In order for server 1 to be idle at time $\bar{\tau}$ (so that $d\hat{L}_1^n(\bar{\tau}) > 0$), it is necessary that $\hat{X}_1^n(\bar{\tau}) = 0$. This means that in order to also have $\hat{W}_{\bar{\tau}}^n \geq c_3 \hat{\Theta}^n$, we need to have $\hat{X}_2^n(\bar{\tau}) \geq 2\hat{\Theta}^n$, or

$$\hat{X}_2^n(\bar{\tau}) - \hat{X}_2^n(\bar{\sigma}) \geq 2\hat{\Theta}^n - n^{-1/2} \geq \hat{\Theta}^n. \quad (\text{A.7})$$

We can now give the balance equation for $\hat{X}_2^n$ between $\bar{\sigma}$ and $\bar{\tau}$ under the $\mathbf{T}_2$ rule. Under the $\mathbf{T}_2$ rule, server 2 prioritizes class 2 for the whole period because the number of HPC jobs remains below $\hat{\Theta}^n$. Because $\lambda_1 > \mu_{11}$ we obtain $\lambda_2 < \mu_{22}$ from (3.1). By the same reasoning as (A.3), there exists $c_2 > 0$ such that

$$\hat{X}_2^n[\bar{\sigma}, \bar{\tau}] \leq \hat{f}_2^n[\bar{\sigma}, \bar{\tau}] - c_2 \sqrt{n}(\bar{\tau} - \bar{\sigma}) + \sqrt{n}e_{\max}^n \sum_k \mu_{2k}. \quad (\text{A.8})$$

Combining (A.7) and (A.8), we obtain (on $\{\rho \geq \bar{\tau}\}$)

$$\hat{f}_2^n[\bar{\sigma}, \bar{\tau}] - c_2 \sqrt{n}(\bar{\tau} - \bar{\sigma}) + \sqrt{n}e_{\max}^n \sum_k \mu_{2k} \geq \hat{\Theta}^n.$$

As in the proof of Proposition 5.4, we distinguish between two cases: $\bar{\tau} - \bar{\sigma}$ smaller or larger than $n^{-\nu_1}$. On $\bar{\tau} - \bar{\sigma} \leq n^{-\nu_1}$, $\hat{f}_2^n[\bar{\sigma}, \bar{\tau}] \leq w_{t_0}(\hat{f}_2^n, n^{-\nu_1})$, so that

$$\begin{aligned} & \mathbb{P} \left(\hat{f}_2^n[\bar{\sigma}, \bar{\tau}] - c_2 \sqrt{n}(\bar{\tau} - \bar{\sigma}) + \sqrt{n}e_{\max}^n \sum_k \mu_{2k} \geq \hat{\Theta}^n, \bar{\tau} - \bar{\sigma} \leq n^{-\nu_1} \right) \\ & \leq \mathbb{P} \left(w_{t_0}(\hat{f}_2^n, n^{-\nu_1}) \geq \hat{\Theta}^n / 2 \right) + \mathbb{P} \left(\sqrt{n}e_{\max}^n \sum_k \mu_{2k} \geq \hat{\Theta}^n / 2 \right). \end{aligned}$$

On $\bar{\tau} - \bar{\sigma} > n^{-\nu_1}$, $\hat{f}_2^n[\bar{\sigma}, \bar{\tau}] \leq 2\|\hat{f}_2^n\|_{t_0}$ and $c_2 \sqrt{n}(\bar{\tau} - \bar{\sigma}) \geq c_2 n^{\frac{1}{2}-\nu_1}$, so that

$$\begin{aligned} & \mathbb{P} \left(\hat{f}_2^n[\bar{\sigma}, \bar{\tau}] - c_2 \sqrt{n}(\bar{\tau} - \bar{\sigma}) + \sqrt{n}e_{\max}^n \sum_k \mu_{2k} \geq \hat{\Theta}^n, \bar{\tau} - \bar{\sigma} > n^{-\nu_1} \right) \\ & \leq \mathbb{P} \left(2\|\hat{f}_2^n\|_{t_0} \geq c_2 n^{\frac{1}{2}-\nu_1} / 2 \right) + \mathbb{P} \left(\sqrt{n}e_{\max}^n \sum_k \mu_{2k} \geq c_2 n^{\frac{1}{2}-\nu_1} / 2 \right). \end{aligned}$$

To summarize,

$$\begin{aligned} \mathbb{P}(\rho \geq \bar{\tau}) & \leq \mathbb{P}(\tau_c^n \leq t_0) + \mathbb{P}(\tau_c^n \geq t_0, \bar{\tau} < \rho) + \mathbb{P}(\tau_c^n \geq t_0, \bar{\tau} \geq \rho, \bar{\tau} \leq \rho, \bar{\tau} - \bar{\sigma} \leq n^{-\nu_1}) + \mathbb{P}(\tau_c^n \geq t_0, \bar{\tau} \geq \rho, \bar{\tau} \leq \rho, \bar{\tau} - \bar{\sigma} > n^{-\nu_1}) \\ & \leq \mathbb{P}(\tau_c^n \leq t_0) + \mathbb{P}(\tau_c^n \geq t_0, \bar{\tau} < \rho) + \mathbb{P} \left(w_{t_0}(\hat{f}_2^n, n^{-\nu_1}) \geq \hat{\Theta}^n / 2 \right) + \mathbb{P} \left(\sqrt{n}e_{\max}^n \sum_k \mu_{2k} \geq \hat{\Theta}^n / 2 \right) \\ & \quad + \mathbb{P} \left(2\|\hat{f}_2^n\|_{t_0} \geq c_2 n^{\frac{1}{2}-\nu_1} / 2 \right) + \mathbb{P} \left(\sqrt{n}e_{\max}^n \sum_k \mu_{2k} \geq c_2 n^{\frac{1}{2}-\nu_1} / 2 \right). \end{aligned}$$

All terms converge to zero. The first two have already been treated in Proposition 5.4 and (A.4), respectively. The third term converges by Lemma 5.7, the fourth and sixth by Lemma 5.9, and the fifth by tightness of $\hat{f}_2^n$.

This completes the proof of (5.26) and Lemma 5.12 in all of the relevant cases. $\square$

Proof of Lemma 5.13.

Case 1(b), T_2 policy:

We begin the proof by a full analysis in the case where a T_2 policy is used and then explain how to deal with the other cases. In this case, with the active mode in canonical form, $p = 1$. Recall the definition of the random time τ :

$$\tau = \inf\{t > \rho : \hat{X}_1^n(t) = 2n^{-1/2}, \text{ and } \hat{X}_2^n(t) \geq \hat{\Theta}^n\},$$

and

$$\sigma = \sup\left\{t \leq \tau^n : \hat{X}_1^n(t) \geq \hat{\Theta}^n, \text{ or } \hat{X}_2^n(t) \leq \frac{\hat{\Theta}^n}{2}\right\}.$$

For any $t \in [\sigma, \tau]$,

- $\hat{X}_1^n(t) < \hat{\Theta}^n$,
- and $\hat{X}_2^n(t) > \frac{\hat{\Theta}^n}{2}$.

In this case, only server 1 processes class 1 jobs and does so at most at full rate. Indeed, only server 1 gives priority to class 1 jobs, and the other server is busy with class 2 jobs present in the system. Similarly, only server 2 processes class 2 jobs and does so at full rate because there are always class 2 jobs and the number of HPC jobs is below Θ^n between σ and τ . Similarly to Lemma 5.10, $e_{\max}^n$ is such that for any $s, t \in [\sigma, \tau]$, $s \leq t$, including a service that could start before σ at the wrong activity (server 2 in this case),

$$\int_s^t \left(\lambda_1 - \sum_k \mu_{1k} \Xi_{1k}^n(s) \right) ds \geq (\lambda_1 - \mu_{11})(t - s) - \mu_{12} e_{\max}^n,$$

and excluding a service of a HPC job by server 2 starting before σ ,

$$\int_s^t \left(\lambda_2 - \sum_k \mu_{2k} \Xi_{2k}^n(s) \right) ds \leq (\lambda_2 - \mu_{22})(t - s) + \mu_{22} e_{\max}^n.$$

With the active mode in canonical form, we have $\lambda_1 > \mu_{11}$. By (3.1), this means $\lambda_2 < \mu_{22}$. Thus, $\lambda_1 - \mu_{11} > 0$ and $\lambda_2 - \mu_{22} < 0$.

There are two possibilities for the state of the system at time σ : we have either

- $\hat{X}_1^n(\sigma) = \hat{\Theta}^n - n^{-1/2}$,
- or $\hat{X}_2^n(\sigma) \in \left[\frac{\hat{\Theta}^n}{2} + \frac{1}{2}n^{-1/2}, \frac{\hat{\Theta}^n}{2} + n^{-1/2} \right]$.

We will decompose the event $\Omega^n := \{\tau \leq t_0\}$ in two events:

$$\Omega^n = \Omega_1^n \cup \Omega_2^n,$$

with

$$\Omega_1^n := \Omega^n \cap \{\hat{X}_1^n(\sigma) = \hat{\Theta}^n - n^{-1/2}\},$$

and

$$\Omega_2^n := \Omega^n \cap \left\{ \hat{X}_2^n(\sigma) \in \left[\frac{\hat{\Theta}^n}{2} + \frac{1}{2}n^{-1/2}, \frac{\hat{\Theta}^n}{2} + n^{-1/2} \right] \right\}.$$

To lighten notation, introduce also

$$\tilde{\Omega}^n := \left\{ \tau \leq t_0, \sup_{t \in [\sigma, \tau]} \hat{X}_1^n(t) < \hat{\Theta}^n, \inf_{t \in [\sigma, \tau]} \hat{X}_2^n(t) > \frac{\hat{\Theta}^n}{2} \right\}.$$

Similarly to Proposition 5.4, let $\nu_2 = 1/2 - \bar{a} \leq 1/4$ and $\nu_1 \in (\nu_2, 1/2)$ so that $\hat{\Theta}^n = n^{-1/2} \lceil n^{1/2-\nu_2} \rceil \geq n^{-\nu_2}$. Then,

$$\begin{aligned} \mathbb{P}(\Omega_1^n) &= \mathbb{P}(\hat{X}_1^n(\sigma) = \hat{\Theta}^n - n^{-1/2}, \hat{X}_1^n(\tau) \leq 2n^{-1/2}, \tilde{\Omega}^n) \\ &\leq \mathbb{P}\left(\hat{X}_1^n(\tau) - \hat{X}_1^n(\sigma) \leq -\frac{\hat{\Theta}^n}{2}, \tilde{\Omega}^n\right) \\ &= \mathbb{P}\left(\hat{X}_1^n(\tau) - \hat{X}_1^n(\sigma) \leq -\frac{\hat{\Theta}^n}{2}, \tilde{\Omega}^n, \tau - \sigma \leq n^{-\nu_1}\right) \\ &\quad + \mathbb{P}\left(\hat{X}_1^n(\tau) - \hat{X}_1^n(\sigma) \leq -\frac{\hat{\Theta}^n}{2}, \tilde{\Omega}^n, \tau - \sigma > n^{-\nu_1}\right). \end{aligned}$$

On $\tilde{\Omega}^n$, we also have

$$\begin{aligned} \hat{X}_1^n(\tau) - \hat{X}_1^n(\sigma) &= \hat{f}_1^n(\tau) - \hat{f}_1^n(\sigma) + \sqrt{n} \int_{\sigma}^{\tau} \left(\lambda_1 - \sum_k \mu_{1k} \Xi_{1k}^n(t) \right) dt \\ &\geq \hat{f}_1^n(\tau) - \hat{f}_1^n(\sigma) + (\lambda_1 - \mu_{11})\sqrt{n}(\tau - \sigma) - \mu_{12}e_{\max}^n \sqrt{n}. \end{aligned}$$

On the event $\{\tilde{\Omega}^n, \tau - \sigma > n^{-\nu_1}\}$, the reasoning is very similar to the proof of Proposition 5.4:

$$\begin{aligned} &\mathbb{P}\left(\hat{X}_1^n(\tau) - \hat{X}_1^n(\sigma) \leq -\frac{\hat{\Theta}^n}{2}, \tilde{\Omega}^n, \tau - \sigma > n^{-\nu_1}\right) \\ &\leq \mathbb{P}\left(2 \sup_{t \leq t_0} |\hat{f}_1^n(t)| + \mu_{12}\sqrt{n}e_{\max}^n \geq (\lambda_1 - \mu_{11})\sqrt{n}n^{-\nu_1}/2\right). \end{aligned} \quad (\text{A.9})$$

The last probability goes to zero by tightness of $\hat{f}_1^n$, $\nu_1 < 1/2$ and Lemma 5.9. When $\tau - \sigma \leq n^{-\nu_1}$ the situation is again the same as in the proof of Proposition 5.4:

$$\begin{aligned} &\mathbb{P}\left(\hat{X}_1^n(\tau) - \hat{X}_1^n(\sigma) \leq -\frac{\hat{\Theta}^n}{2}, \tilde{\Omega}^n, \tau - \sigma \leq n^{-\nu_1}\right) \\ &\leq \mathbb{P}\left(w_{t_0}(\hat{f}_1^n, n^{-\nu_1}) + \mu_{12}e_{\max}^n \sqrt{n} \geq \frac{n^{-\nu_2}}{2}\right). \end{aligned} \quad (\text{A.10})$$

The right-hand side also converges to zero by Lemma 5.7, Remark 5.8, and Lemma 5.9, similarly to the proof of Proposition 5.4.

The second event Ω_2^n is treated similarly:

$$\begin{aligned} \mathbb{P}(\Omega_2^n) &\leq \mathbb{P}\left(\hat{X}_2^n(\sigma) \in \left[\frac{\hat{\Theta}^n}{2} + \frac{1}{2}n^{-1/2}, \frac{\hat{\Theta}^n}{2} + n^{-1/2}\right], \hat{X}_2^n(\tau) \geq \hat{\Theta}^n, \tilde{\Omega}^n\right) \\ &\leq \mathbb{P}\left(\hat{X}_2^n(\tau) - \hat{X}_2^n(\sigma) \geq \frac{\hat{\Theta}^n}{3}, \tilde{\Omega}^n\right) \\ &= \mathbb{P}\left(\hat{X}_2^n(\tau) - \hat{X}_2^n(\sigma) \geq \frac{\hat{\Theta}^n}{3}, \tilde{\Omega}^n, \tau - \sigma \leq n^{-\nu_1}\right) \\ &\quad + \mathbb{P}\left(\hat{X}_2^n(\tau) - \hat{X}_2^n(\sigma) \geq \frac{\hat{\Theta}^n}{3}, \tilde{\Omega}^n, \tau - \sigma > n^{-\nu_1}\right). \end{aligned}$$

On this event, we also have

$$\begin{aligned} \hat{X}_2^n(\tau) - \hat{X}_2^n(\sigma) &= \hat{f}_2^n(\tau) - \hat{f}_2^n(\sigma) + \sqrt{n} \int_{\sigma}^{\tau} \left(\lambda_2 - \sum_k \mu_{2k} \Xi_{2k}^n(t) \right) dt \\ &\leq \hat{f}_2^n(\tau) - \hat{f}_2^n(\sigma) + (\lambda_2 - \mu_{22})\sqrt{n}(\tau - \sigma) + \mu_{22}e_{\max}^n \sqrt{n}. \end{aligned}$$

On the event $\{\tau - \sigma > n^{-\nu_1}\}$, similar to the treatment of Ω_1^n :

$$\begin{aligned} & \mathbb{P}\left(\hat{X}_2^n(\tau) - \hat{X}_2^n(\sigma) \geq \frac{\hat{\Theta}^n}{3}, \tilde{\Omega}^n, \tau - \sigma > n^{-\nu_1}\right) \\ & \leq \mathbb{P}\left(2 \sup_{t \leq t_0} |\hat{f}_2^n(t)| + \mu_{22} \sqrt{n} e_{\max}^n \geq -(\lambda_2 - \mu_{22}) \sqrt{n} n^{-\nu_1} / 3\right). \end{aligned} \quad (\text{A.11})$$

The last probability goes to zero by tightness of $\hat{f}_2^n$, $\nu_1 < 1/2$ and Lemma 5.9. When $\tau - \sigma \leq n^{-\nu_1}$, the situation is also the same as for Ω_1^n :

$$\begin{aligned} & \mathbb{P}\left(\hat{X}_2^n(\tau) - \hat{X}_2^n(\sigma) \geq \frac{\hat{\Theta}^n}{3}, \tilde{\Omega}^n, \tau - \sigma \leq n^{-\nu_1}\right) \\ & \leq \mathbb{P}\left(w_{t_0}(\hat{f}_2^n, n^{-\nu_1}) + \mu_{22} e_{\max}^n \sqrt{n} \geq \frac{n^{-\nu_2}}{3}\right). \end{aligned} \quad (\text{A.12})$$

The right-hand side also converges to zero by Lemma 5.7, similarly to the Ω_1^n case. This concludes the proof in case 1(b). We now present the changes required to adapt the proof to the other cases described in the lemma:

Case 2(b), T_2 T_2 policy:

This case involves switching between two T_2 rules: $i_2(\xi^L) = i_2(\xi^H) = p$. Because ξ^L is in canonical form, $p = 1$. In this case, mode switching only occurs at a service completion at the single-activity server that is dedicated to HPC jobs. There could be either type of job being served at either server immediately before time σ . We will see that after those jobs exit the system, there is at least one server processing the LPC and at most one processing the HPC. If the first job to finish is from the dual-activity server, the service of a low-priority job starts because there are fewer than Θ^n HPC jobs. This server will continue to take LPC jobs until the minimum between the next mode switching time and τ . If the first job to finish at or after σ is at the single-activity server, either the current mode switches or the service of an HPC job starts. If there is no mode switch, then this server will continue to take HPC jobs until the minimum between the next mode switching time and τ . When the dual-activity server completes its job, it will, as noted before, serve LPC jobs. If there is a mode switch, then the formerly single-activity server becomes dual activity, and (because there are fewer than Θ^n HPC jobs) begins service on an LPC job. When the formerly dual-activity server completes its job, it becomes single activity and serves HPC. This continues until the minimum between the next mode switching time and τ .

After the two jobs present at σ , there cannot be a time where both servers are occupied with HPC jobs. This is because the HPC is only processed by the server dedicated to it, and there can be no residual service of an HPC job at the single activity server whenever the current mode switches. Similarly, after both initial jobs have been processed, there is always at least one server occupied with LPC jobs. This is because one server gives priority to the LPC, and the service of a LPC job starts whenever the current mode switches.

Keeping the definition of τ, σ , just as in the single-mode case, for any $t, s \in [\sigma, \tau]$, including/excluding a service that could start before σ at the wrong activity, there is always at most one server processing HPC jobs and at least one server processing LPC jobs between σ and τ . Thus,

$$\int_s^t \left(\lambda_1 - \sum_k \mu_{1k} \Xi_{1k}^n(s) \right) ds \geq \left(\lambda_1 - \max_k \mu_{1k} \right) (t - s) - \sum_k \mu_{1k} e_{\max}^n,$$

and

$$\int_s^t \left(\lambda_2 - \sum_k \mu_{2k} \Xi_{2k}^n(s) \right) ds \leq \left(\lambda_2 - \min_k \mu_{2k} \right) (t - s) + \sum_k \mu_{2k} e_{\max}^n.$$

With one mode in canonical form we have $\frac{\lambda_1}{\alpha_1} > \beta_1$. In addition, by Lemma 3.3, in case 2(b), (3.6) holds, and, thus, $\frac{\lambda_1}{\alpha_1} > \max_k \beta_k$, which also means $\frac{\lambda_2}{\alpha_2} < \min_k \beta_k$ by (3.1).

The rest of the proof is the same as in the single-mode case, obtaining (A.9), (A.10), (A.11), and (A.12). This concludes the proof in case 2(b).

Case 2(c), $T_1 T_2$ policy:

With ξ^L in canonical form, we have $p = 2$, so τ and σ are defined as

$$\tau = \inf\{t > \rho : \hat{X}_2^n(t) = 2n^{-1/2}, \text{ and } \hat{X}_1^n(t) \geq \hat{\Theta}^n\},$$

and

$$\sigma = \sup \left\{ t \leq \tau^n : \hat{X}_2^n(t) \geq \hat{\Theta}^n, \text{ or } \hat{X}_1^n(t) \leq \frac{\hat{\Theta}^n}{2} \right\}.$$

In this case, in the upper-workload mode, only the single-activity server is allowed to begin service of the HPC between σ and τ , whereas in the lower-workload mode, neither server can begin service of HPC between σ and τ . Between those times, there are always low-priority class jobs to serve, and the number of HPC jobs stays below Θ^n . The most service the HPC can get between σ and τ occurs if there are no switches and the current mode is always upper workload. In this mode, the single-activity server (server 1) is dedicated to service of the HPC. Hence, including a service that could start before σ at the wrong activity,

$$\int_s^t \left(\lambda_2 - \sum_k \mu_{2k} \Xi_{2k}^n(s) \right) ds \geq (\lambda_2 - \mu_{21})(t - s) - \mu_{21} e_{\max}^n.$$

Between σ and τ , server 2 gives priority to class 1 regardless of switches. Excluding a service of a HPC job that could have started before σ ,

$$\int_s^t \left(\lambda_1 - \sum_k \mu_{1k} \Xi_{1k}^n(s) \right) ds \leq (\lambda_1 - \mu_{12})(t - s) - \mu_{12} e_{\max}^n.$$

Because we have one mode in canonical form, $\frac{\lambda_1}{\alpha_1} > \beta_1$. In addition, by Lemma 3.3 in case 2(c), (3.5) holds, which means that $\frac{\lambda_1}{\alpha_1} < \beta_2$. Finally, $\frac{\lambda_2}{\alpha_2} > \beta_1$ by the previous observation and (3.1).

The rest of the proof is the same as in the single-mode case, obtaining (A.9), (A.10), (A.11), and (A.12).

Case 2(d), $T_2 T_1$ policy:

This case is handled the same way as 2(c), by interchanging the roles of upper workload and lower workload mode. $\square$

Proof of Lemma 5.11 (Continued from Section 5.2.5).

• **Case 1(b), T_2 policy:**

In this case, we already know by Lemma 5.12 that $\mathbb{P} \left(\bar{R}_\rho^n \geq \frac{\delta}{2} \right) \rightarrow 0$. We need to show the same thing for the time after ρ :

$$\mathbb{P} \left(\int_\rho^{t_0} \mathbb{1}_{W_t^n \geq c_3 \Theta^n} d\hat{L}_t^n \geq \frac{\delta}{2} \right) \rightarrow 0.$$

First, by putting ξ^A in canonical form, $p = 1$. When $X_1^n(t)$ is above the threshold, both servers can serve class 1 jobs (high-priority class), so almost surely,

$$\int_\rho^{t_0} \mathbb{1}_{\hat{X}_1^n(t) \geq \hat{\Theta}^n} d\hat{L}_t^n = 0. \quad (\text{A.13})$$

Similarly,

$$\int_\rho^{t_0} \mathbb{1}_{\hat{X}_2^n(t) \geq 2n^{-1/2}} d\hat{L}_t^n = 0, \quad (\text{A.14})$$

and

$$\int_\rho^{t_0} \mathbb{1}_{2n^{-1/2} \leq \hat{X}_1^n(t) \leq \hat{\Theta}^n} d\hat{L}_t^n = 0. \quad (\text{A.15})$$

Notice now that because of the three identities,

$$\begin{aligned} \mathbb{P} \left(\int_\rho^{t_0} \mathbb{1}_{W_t^n \geq c_3 \Theta^n} d\hat{L}_t^n \geq \frac{\delta}{2} \right) &\leq \mathbb{P} \left(\int_\rho^{t_0} \left(\mathbb{1}_{\hat{X}_1^n(t) \geq \hat{\Theta}^n} + \mathbb{1}_{\hat{X}_1^n(t) \leq \hat{\Theta}^n, \hat{X}_2^n(t) \geq \hat{\Theta}^n} \right) d\hat{L}_t^n \geq \frac{\delta}{2} \right) \\ &= \mathbb{P} \left(\beta_1 \int_\rho^{t_0} \mathbb{1}_{\hat{X}_1^n(t) \leq \hat{\Theta}^n, \hat{X}_2^n(t) \geq \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2} \right) \\ &= \mathbb{P} \left(\beta_1 \int_\rho^{t_0} \mathbb{1}_{\hat{X}_1^n(t) \leq 2n^{-1/2}, \hat{X}_2^n(t) \geq \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2} \right). \end{aligned}$$

By Lemma 5.13,

$$\mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{\hat{X}_1^n(t) \leq 2n^{-1/2}, \hat{X}_2^n \geq \hat{\Theta}^n} d\hat{I}_1^n(t) \geq \frac{\delta}{2}\right) \leq \mathbb{P}(\tau_r^n \leq t_0) \rightarrow 0. \quad (\text{A.16})$$

• **Case 2(a), PP policy:**

In this case, $p = 2$. The reasoning in case 1(a) is still valid when switching between two P rules because (5.28) and (5.29) still hold.

• **Case 2(b), T_2T_2 policy:**

The result in this case has a proof very similar to case 1(b) because Lemmas 5.12 and 5.13 are still valid in this case. We already know by Lemma 5.12 that

$$\mathbb{P}\left(\int_0^{\rho} \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2}\right) \rightarrow 0.$$

We need to show the same thing for the time after ρ :

$$\mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2}\right) \rightarrow 0.$$

First, by putting ξ^L in canonical form $p = 1$.

The idea is to split the integrals between the times $\text{MODE}(t) = \xi^L$ and the times $\text{MODE}(t) = \xi^H$. In this case, we still have (A.13), but this time (A.14) and (A.15) only hold when $\text{MODE}(t) = \xi^L$. In the other case, we have

$$\begin{aligned} \int_{\rho}^{t_0} \mathbb{1}_{\hat{X}_2^n(t) \geq 2n^{-1/2}, \text{MODE}(t) = \xi^H} d\hat{I}_1^n(t) &= 0, \\ \int_{\rho}^{t_0} \mathbb{1}_{2n^{-1/2} \leq \hat{X}_1^n(t) \leq \hat{\Theta}^n, \text{MODE}(t) = \xi^H} d\hat{I}_2^n(t) &= 0. \end{aligned}$$

We obtain

$$\int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t) = \xi^L, \hat{W}_t^n \geq c_3 \hat{\Theta}^n} d\hat{L}_t^n \leq \beta_1 \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t) = \xi^L, \hat{X}_1^n(t) \leq 2n^{-1/2}, \hat{X}_2^n(t) \geq \hat{\Theta}^n} d\hat{I}_1^n(t),$$

and

$$\int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t) = \xi^H, \hat{W}_t^n \geq c_3 \hat{\Theta}^n} d\hat{L}_t^n \leq \beta_2 \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t) = \xi^H, \hat{X}_1^n(t) \leq 2n^{-1/2}, \hat{X}_2^n(t) \geq \hat{\Theta}^n} d\hat{I}_2^n(t).$$

Finally, we obtain the result, because

$$\mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{\hat{X}_1^n(t) \leq 2n^{-1/2}, \hat{X}_2^n(t) \geq \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2}\right) \leq \mathbb{P}(\tau_r^n \leq t_0) \rightarrow 0.$$

• **Case 2(c)(d), T_1T_2/T_2T_1 policy:**

We give the proof in case 2(d), but the reasoning is similar in case 2(c) by interchanging the role of ξ^L and ξ^H . In this case, $p = 1$ with ξ^L in canonical form. The idea is similar to the 2(b) case. We already know by Lemma 5.12 that

$$\mathbb{P}\left(\int_0^{\rho} \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2}\right) \rightarrow 0.$$

We need to show the same thing for the time after ρ :

$$\mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{\hat{W}_t^n \geq c_3 \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{2}\right) \rightarrow 0.$$

We will split the integral using $\mathbb{1}_{\text{MODE}(t)=\xi^L} + \mathbb{1}_{\text{MODE}(t)=\xi^H}$. We use a T_2 rule when W^n is small and a T_1 rule when W^n is large, but that distinction is not important here. In terms of almost sure nonidling properties, we have

$$\begin{aligned} \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^L, \hat{X}_1^n(t) \geq \hat{\Theta}^n} d\hat{L}_t^n &= 0, \quad \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^L, \hat{X}_2^n(t) \geq 2n^{-1/2}} d\hat{L}_t^n(t) = 0, \\ \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^L, 2n^{-1/2} \leq \hat{X}_1^n(t) \leq \hat{\Theta}^n} d\hat{L}_t^n(t) &= 0, \quad \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^H, \hat{X}_2^n(t) \geq 2n^{-1/2}} d\hat{L}_t^n = 0. \end{aligned}$$

From these almost sure identities, we obtain

$$\begin{aligned} \mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{W_t^n \geq c_3 \Theta^n} d\hat{L}_t^n \geq \frac{\delta}{2}\right) &\leq \mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^L, W_t^n \geq c_3 \Theta^n} d\hat{L}_t^n \geq \frac{\delta}{4}\right) \\ &\quad + \mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^H, W_t^n \geq c_3 \Theta^n} d\hat{L}_t^n \geq \frac{\delta}{4}\right) \\ &\leq \mathbb{P}\left(\beta_1 \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^L, \hat{X}_1^n(t) \leq 2n^{-1/2}, \hat{X}_2^n(t) \geq \hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{4}\right) \\ &\quad + \mathbb{P}\left(\int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^H, \hat{X}_2^n(t) \leq 2n^{-1/2}, \hat{X}_1^n(t) \geq 2\hat{\Theta}^n} d\hat{L}_t^n \geq \frac{\delta}{4}\right) \\ &\leq \mathbb{P}(\tau_r^n \leq t_0) + \mathbb{P}(\tau_c^n \leq t_0). \end{aligned}$$

Both probabilities go to zero (by Lemma 5.13 for the first and Proposition 5.4 for the second), so the result is proved in this case as well.

As mentioned, the proof is the same in case 2(c): this time, the policy uses a T_1 rule when W^n is small and T_2 rule when W^n is large. Keeping ξ^L in canonical form, $p = 2$. In addition, in terms of almost sure nonidling properties, we have

$$\begin{aligned} \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^H, \hat{X}_2^n(t) \geq \hat{\Theta}^n} d\hat{L}_t^n &= 0, \quad \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^H, \hat{X}_1^n(t) \geq 2n^{-1/2}} d\hat{L}_t^n(t) = 0, \\ \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^H, 2n^{-1/2} \leq \hat{X}_2^n(t) \leq \hat{\Theta}^n} d\hat{L}_t^n(t) &= 0, \quad \int_{\rho}^{t_0} \mathbb{1}_{\text{MODE}(t)=\xi^L, \hat{X}_1^n(t) \geq 2n^{-1/2}} d\hat{L}_t^n = 0. \end{aligned}$$

We can obtain the result using the same decomposition. $\square$

Appendix B. Solution of the HJB Equation and Free Boundary Point

In this appendix, we present the expression found in Sheng (1981, section 3, case 1) for the solution to the HJB equation. (This solution is a slightly corrected version of the solution presented in Sheng 1978, section 5.3.) It includes an equation that uniquely characterizes the free boundary point (or switching point) z^* in the dual-mode case. It is assumed in Sheng (1981), without loss of generality, that $b_1 \geq b_2$. We assume further, for simplicity (and, again, without loss of generality), that if $b_1 = b_2$, then $\sigma_1 \geq \sigma_2$. Note that, with these indexing assumptions, $m = 2$ in whichever of the complementary Conditions (2.20) or (2.21) that holds. Throughout this section, denote the unique classical solution to (2.18)–(2.19) by $u(x)$, $x \in \mathbb{R}_+$, and let x serve as the initial condition for the WCP, which elsewhere in this paper is denoted by z . Let

$$\beta = \frac{b_1 + \sqrt{b_1^2 + 2\gamma\sigma_1^2}}{\sigma_1^2}, \quad \rho = \frac{b_2 + \sqrt{b_2^2 + 2\gamma\sigma_2^2}}{\sigma_2^2}, \quad v = \frac{b_1 - \sqrt{b_1^2 + 2\gamma\sigma_1^2}}{\sigma_1^2}.$$

Theorem B.1 (Sheng 1981, Section 3, Case 1). *Under (2.20), $u(x) = \frac{x}{\gamma} + \frac{b_2}{\gamma^2} + \frac{1}{\gamma\rho} e^{-\rho x}$, $x \in \mathbb{R}_+$.*

Next, consider Condition (2.21). Because b_1 and b_2 are distinct in this case, we have $b_1 > b_2$. The Policy (4.1) from Section 4 corresponding to switching at z is given in the present notation by $\bar{\xi}_z(x) = \xi^{*,1} \mathbb{1}_{x \leq z} + \xi^{*,2} \mathbb{1}_{x > z}$. Let $\mathfrak{S}_z^{(2)}$ be the admissible control system from Lemma 4.1.2, with a generic switching point z in place of the specific z^* . Let the corresponding expression

$J_{\text{WCP}}(x, \bar{\xi}_z^{(2)})$, which is nothing but the cost associated with the switching policy $\bar{\xi}_z$, be denoted by $J(x, \bar{\xi}^z)$. Following is an expression for this cost. For $z > 0$, let

$$\begin{aligned} A(z) &= \frac{\nu\beta(e^{(\nu-\beta)z} + e^{-(\nu-\beta)z} - 2)}{\nu\beta(e^{(\nu-\beta)z} + e^{-(\nu-\beta)z} - 2) + \rho(e^{-\nu z} - e^{-\beta z})(\nu e^{\nu z} - \beta e^{\beta z})}, \\ C(z) &= \frac{(\nu - \beta)(e^{-\nu z} - e^{-\beta z})}{\nu\beta(e^{(\nu-\beta)z} + e^{-(\nu-\beta)z} - 2) + \rho(e^{-\nu z} - e^{-\beta z})(\nu e^{\nu z} - \beta e^{\beta z})}, \\ D(z) &= \frac{(\beta - \nu)\rho(e^{\beta z} - e^{\nu z})}{\nu\beta(e^{(\nu-\beta)z} + e^{-(\nu-\beta)z} - 2) + \rho(e^{-\nu z} - e^{-\beta z})(\nu e^{\nu z} - \beta e^{\beta z})}, \\ F(z) &= \frac{e^{-\nu z} - e^{-\beta z} + (\beta - \nu)C(z)}{\nu e^{-\beta z} - \beta e^{-\nu z}}. \end{aligned}$$

($F(\cdot)$ is not to be confused with the process F defined in the body of the paper). Then, for $x \leq z$,

$$\begin{aligned} J(x, \bar{\xi}^z) &= \frac{x}{\gamma} + \frac{b_1[(e^{-\nu z} - e^{-\beta z}) + (A(z) - 1)(e^{-\nu x} - e^{-\beta x}) + D(z)(e^{-\nu z - \beta x} - e^{-\beta z - \nu x})]}{\gamma^2(e^{-\nu z} - e^{-\beta z})} \\ &\quad + \frac{b_2[(1 - A(z))(e^{-\nu x} - e^{-\beta x}) - D(z)(e^{-\nu z - \beta x} - e^{-\beta z - \nu x})]}{\gamma^2(e^{-\nu z} - e^{-\beta z})} \\ &\quad + \frac{C(z)(e^{-\nu x} - e^{-\beta x}) - F(z)(e^{-\nu z - \beta x} - e^{-\beta z - \nu x})}{\gamma(e^{-\nu z} - e^{-\beta z})}, \end{aligned} \quad (\text{B.1})$$

and for $x \geq z$,

$$J(x, \bar{\xi}^z) = \frac{x}{\gamma} + \frac{b_1 A(z)}{\gamma^2 e^{\rho(x-z)}} + \frac{b_2}{\gamma^2} (1 - e^{-\rho(x-z)} A(z)) + \frac{C(z)}{\gamma e^{\rho(x-z)}}. \quad (\text{B.2})$$

It is here where the principle of smooth fit is applied. For the cost to be C^2 (in x), it must satisfy $J''(z-, \bar{\xi}^z) = J''(z+, \bar{\xi}^z)$. Using the Expressions (B.1) and (B.2), this condition can be translated to the following equation:

$$\begin{aligned} 0 &= \left(\frac{b_1 - b_2}{\gamma^2} \right) \left[\frac{(A(z) - 1)(\nu^2 e^{-\nu z} - \beta^2 e^{-\beta z})}{e^{-\nu z} - e^{-\beta z}} - \rho^2 A(z) + \frac{(\beta^2 - \nu^2) D(z) e^{-(\nu+\beta)z}}{e^{-\nu z} - e^{-\beta z}} \right] \\ &\quad + \frac{1}{\gamma} \left[\frac{C(z)(\nu^2 e^{-\nu z} - \beta^2 e^{-\beta z})}{e^{-\nu z} - e^{-\beta z}} - \rho^2 C(z) - \frac{(\beta^2 - \nu^2) F(z) e^{-(\nu+\beta)z}}{e^{-\nu z} - e^{-\beta z}} \right]. \end{aligned} \quad (\text{B.3})$$

Theorem B.2 (Sheng 1981, Section 3, Case 1). Let (2.21) hold. Then, (B.3) has a unique solution $z^* \in (0, \infty)$. Moreover, $u(x) = J(x, \bar{\xi}^{z^*})$, $x \in \mathbb{R}_+$.

Appendix C. Symmetry Conditions

The following result is related to Remark 4.2.

Lemma C.1.

1. If either (4.2) or (4.3) holds, then $b_1 = b_2$. In particular, (2.20) holds.
2. If (4.4) holds, then $\sigma_1 = \sigma_2$. In particular, (2.20) holds.

Proof.

1. Recall the expressions for b_1 and b_2 ,

$$b_m = b(\xi^{*,m}) = \sum_i \frac{\hat{\lambda}_i - \sum_k \hat{\mu}_{ik}^{\xi^{*,m}}}{\alpha_i}.$$

The difference between b_1 and b_2 is thus the difference between $\gamma_m := \sum_{i,k} \frac{\hat{\mu}_{ik}^{\xi^{*,m}}}{\alpha_i}$. For $\xi^{*,1}$, we distinguish these cases: either $\frac{\lambda_1}{\alpha_1} > \beta_2$

or $\frac{\lambda_1}{\alpha_1} < \beta_2$. For $\xi^{*,2}$, we distinguish these cases: either $\frac{\lambda_1}{\alpha_1} < \beta_1$ or $\frac{\lambda_1}{\alpha_1} > \beta_1$. We will see that in each of the four cases $b_1 = b_2$,

$$\begin{aligned} \frac{\lambda_1}{\alpha_1} > \beta_2 : \quad \gamma_1 &= \frac{1}{\alpha_1} \left[\hat{\mu}_{12} + \hat{\mu}_{11} \left(\frac{\lambda_1}{\alpha_1 \beta_1} - \frac{\beta_2}{\beta_1} \right) \right] + \frac{\hat{\mu}_{21}}{\alpha_2} \left(1 - \frac{\lambda_1}{\alpha_1 \beta_1} + \frac{\beta_2}{\beta_1} \right), \\ \frac{\lambda_1}{\alpha_1} < \beta_2 : \quad \gamma_1 &= \frac{\hat{\mu}_{12} \lambda_1}{\alpha_1^2 \beta_2} + \frac{1}{\alpha_2} \left[\hat{\mu}_{21} + \hat{\mu}_{22} \left(1 - \frac{\lambda_1}{\alpha_1 \beta_2} \right) \right], \\ \frac{\lambda_1}{\alpha_1} < \beta_1 : \quad \gamma_2 &= \frac{\hat{\mu}_{11} \lambda_1}{\alpha_1^2 \beta_1} + \frac{1}{\alpha_2} \left[\hat{\mu}_{22} + \hat{\mu}_{21} \left(1 - \frac{\lambda_1}{\alpha_1 \beta_1} \right) \right], \\ \frac{\lambda_1}{\alpha_1} > \beta_1 : \quad \gamma_2 &= \frac{1}{\alpha_1} \left[\hat{\mu}_{11} + \hat{\mu}_{12} \left(\frac{\lambda_1}{\alpha_1 \beta_2} - \frac{\beta_1}{\beta_2} \right) \right] + \frac{\hat{\mu}_{22}}{\alpha_2} \left(1 - \frac{\lambda_1}{\alpha_1 \beta_2} + \frac{\beta_1}{\beta_2} \right). \end{aligned}$$

We now take the difference for each pair:

$$\begin{aligned} \frac{\lambda_1}{\alpha_1} > \beta_2 \quad \frac{\lambda_1}{\alpha_1} < \beta_1 : \quad b_1 - b_2 &= \frac{\hat{\mu}_{12} - \hat{\mu}_{11} \frac{\beta_2}{\beta_1}}{\alpha_1} - \frac{\hat{\mu}_{22} - \hat{\mu}_{21} \frac{\beta_2}{\beta_1}}{\alpha_2}, \\ \frac{\lambda_1}{\alpha_1} > \beta_2 \quad \frac{\lambda_1}{\alpha_1} > \beta_1 : \quad b_1 - b_2 &= \left(1 - \frac{\lambda_1}{\alpha_1 \beta_1} + \frac{\beta_2}{\beta_1} \right) \left[\frac{1}{\alpha_1} \left(\hat{\mu}_{12} \frac{\beta_1}{\beta_2} - \hat{\mu}_{11} \right) + \frac{1}{\alpha_2} \left(\hat{\mu}_{21} - \hat{\mu}_{22} \frac{\beta_1}{\beta_2} \right) \right], \\ \frac{\lambda_1}{\alpha_1} < \beta_2 \quad \frac{\lambda_1}{\alpha_1} < \beta_1 : \quad b_1 - b_2 &= \frac{\lambda_1}{\alpha_1^2} \left[\frac{\hat{\mu}_{12}}{\beta_2} - \frac{\hat{\mu}_{11}}{\beta_1} \right] + \frac{\lambda_1}{\alpha_1 \alpha_2} \left[\frac{\hat{\mu}_{21}}{\beta_1} - \frac{\hat{\mu}_{22}}{\beta_2} \right], \\ \frac{\lambda_1}{\alpha_1} < \beta_2 \quad \frac{\lambda_1}{\alpha_1} > \beta_1 : \quad b_1 - b_2 &= \frac{\hat{\mu}_{12} \frac{\beta_1}{\beta_2} - \hat{\mu}_{11}}{\alpha_1} + \frac{\hat{\mu}_{21} - \hat{\mu}_{22} \frac{\beta_1}{\beta_2}}{\alpha_2}. \end{aligned}$$

Now, if $\frac{\hat{\mu}_{i1}}{\beta_1} = \frac{\hat{\mu}_{i2}}{\beta_2}$, $i = 1, 2$, we get

$$\hat{\mu}_{11} \frac{\beta_2}{\beta_1} - \hat{\mu}_{12} = \hat{\mu}_{22} - \hat{\mu}_{21} \frac{\beta_2}{\beta_1} = 0,$$

and, consequently, $b_1 = b_2$ as claimed. If $\frac{\hat{\mu}_{ik}}{\alpha_k} = \frac{\hat{\mu}_{jk}}{\alpha_j}$, $k = 1, 2$, it is not hard to see that again the expressions can be rewritten with a different factorization to get $b_1 = b_2$.

2. Under (4.4), denote $C_i = C_{S_{i1}} = C_{S_{i2}}$. Then, for $\xi \in \mathcal{S}_{LP}$,

$$\begin{aligned} \sigma(\xi)^2 &= \sum_i \frac{\sigma_{A,i}^2 + \sum_k \sigma_{S_{ik}}^2 \xi_{ik}}{\alpha_i^2} \\ &= \sum_i \frac{\sigma_{A,i}^2 + \sum_k C_i^2 \mu_{ik} \xi_{ik}}{\alpha_i^2} \\ &= \sum_i \frac{\sigma_{A,i}^2 + C_i^2 \lambda_i}{\alpha_i^2}, \end{aligned}$$

where the last equality follows from (2.7). Hence $\sigma_1 = \sigma_2$ as claimed. $\square$

References

- Atar R, Lev-Ari A (2018) Workload-dependent dynamic priority for the multiclass queue with reneging. *Math. Oper. Res.* 43(2):494–515.
- Atar R, Castiel E, Reiman MI (2024) Parallel server systems under an extended heavy traffic condition: A lower bound. *Ann. Appl. Probab.* 34(1B):1029–1071.
- Bell SL, Williams RJ (2001) Dynamic scheduling of a system with two parallel servers in heavy traffic with resource pooling: Asymptotic optimality of a threshold policy. *Ann. Appl. Probab.* 11(3):608–649.
- Billingsley P (1999) *Convergence of Probability Measures*, 2nd ed., Wiley Series in Probability and Statistics (Wiley-Interscience, New York).
- Ethier SN, Kurtz TG (1986) *Markov Processes, Characterization and Convergence*, Wiley Series in Probability and Mathematical Statistics (John Wiley & Sons Inc., New York).
- Fleming WH, Soner HM (2006) *Controlled Markov Processes and Viscosity Solutions*, Stochastic Modelling and Applied Probability, vol. 25 (Springer Science & Business Media, New York).
- Goldfarb D, Todd MJ (1989) Linear programming. Nemhauser GL, Rinnoy Kan AHG, Todd MJ, eds. *Optimization*, Handbooks in Operations Research and Management Science, vol. 1 (Elsevier Science, Amsterdam), 73–170.
- Harrison JM (1998) Heavy traffic analysis of a system with parallel servers: Asymptotic optimality of discrete-review policies. *Ann. Appl. Probab.* 8(3):822–848.

- Harrison JM, López MJ (1999) Heavy traffic resource pooling in parallel-server systems. *Queueing Systems Theory Appl.* 33(4):339–368.
- Krichagina EV, Taksar MI (1992) Diffusion approximation for GI/G/1 controlled queues. *Queueing Systems* 12:333–368.
- Krylov NV, Liptser R (2002) On diffusion approximation with discontinuous coefficients. *Stochastic Processes Appl.* 102(2):235–264.
- Kurtz TG, Protter P (1991) Weak limit theorems for stochastic integrals and stochastic differential equations. *Ann. Probab.* 19(3):1035–1070.
- Liptser RS, Shiriaev AN (1977) *Statistics of Random Processes I: General Theory*, Stochastic Modelling and Applied Probability, vol. 5 (Springer, New York).
- Revuz D, Yor M (1999) *Continuous Martingales and Brownian Motion*, 3rd ed., Grundlehren der Mathematischen Wissenschaften [Fundamental Principles of Mathematical Sciences], vol. 293 (Springer-Verlag, Berlin).
- Schrijver A (1986) *Theory of Linear and Integer Programming*, Wiley-Interscience Series in Discrete Mathematics (John Wiley & Sons, Ltd., Chichester, UK).
- Sheng D (1978) Some problems in the optimal control of diffusions. PhD thesis, Stanford University, Stanford, CA.
- Sheng D (1981) Two-mode control of reflecting Brownian motion. Working paper, Bell Labs, Holmdel, NJ.
- Stroock DW, Varadhan SS (2006) *Multidimensional Diffusion Processes*, Classics Math, vol. 233 (Springer-Verlag, Berlin, Heidelberg), 1–338.