



Transportation Science

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

A Slot-Scheduling Mechanism at Congested Airports that Incorporates Efficiency, Fairness, and Airline Preferences

Jamie Fairbrother, Konstantinos G. Zografos, Kevin D. Glazebrook

To cite this article:

Jamie Fairbrother, Konstantinos G. Zografos, Kevin D. Glazebrook (2020) A Slot-Scheduling Mechanism at Congested Airports that Incorporates Efficiency, Fairness, and Airline Preferences. *Transportation Science* 54(1):115-138. <https://doi.org/10.1287/trsc.2019.0926>

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 “*Transportation Science*. Copyright © 2019 The Author(s). <https://doi.org/10.1287/trsc.2019.0926>, used under a Creative Commons Attribution License: <https://creativecommons.org/licenses/by/4.0/>.”

Copyright © 2019, 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>

A Slot-Scheduling Mechanism at Congested Airports that Incorporates Efficiency, Fairness, and Airline Preferences

Jamie Fairbrother,^a Konstantinos G. Zografos,^a Kevin D. Glazebrook^a

^a Centre for Transport and Logistics, Department of Management Science, Lancaster University Management School, Lancaster University, Lancaster LA1 4YX, United Kingdom

Contact: j.fairbrother@lancaster.ac.uk,  <https://orcid.org/0000-0001-5134-4397> (JF); k.zografos@lancaster.ac.uk,

 <https://orcid.org/0000-0002-5779-1419> (KGZ); k.glazebrook@lancaster.ac.uk,  <https://orcid.org/0000-0002-5045-0718> (KDG)

Received: August 20, 2018

Revised: March 2, 2019

Accepted: May 8, 2019

Published Online in Articles in Advance:
December 23, 2019

<https://doi.org/10.1287/trsc.2019.0926>

Copyright: © 2019 The Author(s)

Abstract. Congestion is a problem at airports where capacity does not meet demand. At many such airports, airlines must request time slots for the purpose of landing or take off. Given the imbalance between demand and capacity, slot requests cannot always be scheduled as requested. The difference between the requested and allocated time slots is called displacement. Minimization of the total displacement is a key slot-scheduling objective and expresses the efficiency of the slot-scheduling process. Additionally, fairness has been proposed as a slot-scheduling criterion. Fairness relates to the allocation of the total schedule displacement among the various airlines. Single- and multiobjective models have been proposed for slot scheduling. However, currently the literature lacks models that incorporate the preferences of airlines regarding the allocation of displacement to their flights. This paper proposes a two-stage mechanism for the scheduling of slots at congested airports. The proposed mechanism considers efficiency and fairness objectives and incorporates the preferences of airlines in allocating the total displacement associated with the flights of each airline. The first stage of the mechanism constructs a reference schedule that is fair to the participating airlines. In the second stage, the airlines specify how the displacement allocated to them in the reference schedule should be distributed among their requests. The mechanism then adjusts the fair reference schedule to meet as many of these preferences as possible. The development and implementation of the proposed slot-scheduling mechanism is demonstrated using real data from a coordinated airport and simulated displacement preference data. The proposed slot-scheduling mechanism provides useful information to decision makers regarding the equity–efficiency trade-off and enhances the transparency and acceptability of the slot-scheduling outcome.



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 “Transportation Science. Copyright © 2019 The Author(s). <https://doi.org/10.1287/trsc.2019.0926>, used under a Creative Commons Attribution License: <https://creativecommons.org/licenses/by/4.0/>.”

Funding: This work was supported by the Engineering and Physical Sciences Research Council [Grant EP/M020258/1].

Supplemental Material: The online appendix is available at <https://doi.org/10.1287/trsc.2019.0926>.

Keywords: airport slot scheduling • airport capacity management • slot-scheduling mechanism • biobjective slot scheduling • slot-scheduling fairness • integer linear programming

1. Introduction

Because of insufficient capacity to meet demand, many airports around the world suffer congestion-related delays and scheduling infeasibilities. Physical and political constraints mean that the expansion of infrastructure is not possible in the short to medium term, and so congestion must be mitigated through the management of demand.

Various market-based mechanisms such as congestion pricing (Morrison and Winston 2007, Vaze and Barnhart 2012) and auctions (Rassenti, Smith, and Bulfin 1982; Ball, Donohue, and Hoffman 2005) have been proposed in the academic literature. However, the dominant mechanism for managing demand at

congested airports is an administrative scheme called the International Air Transport Association (2017; IATA) Worldwide Slot Guidelines (WSG). This is used throughout the world outside of the United States, and even forms the basis of European Union law on slot allocation (Council of the European Union 1993a, b). Under the IATA WSG, congested airports are designated as *coordinated*, and to use airport infrastructure at a coordinated airport, airlines must request and be allocated time slots for landing or take-off. Although the WSG specify the general procedures, priorities, and constraints that must be followed, they leave some discretion to slot coordinators on how exactly slots are allocated.

Various models for slot allocation have been proposed (Zografos, Salouras, and Madas 2012; Ribeiro et al. 2018; Zografos and Jiang 2019). Although some of these models incorporate fairness, the fairness schemes treat all requests equally, which indirectly penalizes requests made at off-peak times. Moreover, none of these models explicitly take into account the preferences of airlines regarding how their requests should be displaced. These issues are important because a schedule that is fair and takes into account an airline's preferences is much more likely to be acceptable to that airline. The main contribution of this paper is a mechanism that implements the main features of the IATA WSG scheme and also incorporates fairness and the preferences of the airlines. The mechanism works by first constructing a fair reference schedule and then adjusting this according to airline preferences. The proposed fairness scheme improves on previous schemes by taking into account demand for slots at the requested times. In order to provide a detailed description of this paper's contribution, we must first review the IATA WSG and the current literature.

This paper is organized as follows. In the remainder of this section, we review the IATA WSG, the current literature, and describe in detail the contributions of this paper. In Section 2, we provide a summary of the proposed mechanism and the main processes involved. In Section 3, we present a basic slot-allocation model, originally proposed in Zografos, Salouras, and Madas (2012), which is extended in later sections to incorporate fairness and airline preferences. In Section 4, we describe a new fairness scheme. In Section 5, we describe our mechanism for incorporating airline displacement preferences. In Section 6, we explain how the proposed mechanism can be used to take into account the main IATA WSG priority classes. In Section 7, we test our methodology using request data for a coordinated airport and simulated preference data. Finally, in Section 8, we make some concluding remarks and discuss future work.

1.1. IATA Slot-Allocation Process and Priority Classes

In this section, we describe the main features of the IATA slot-allocation process. For a more comprehensive description of the IATA WSG, we refer the reader directly to International Air Transport Association (2017).

Slots at a coordinated airport are allocated for a six-month season, and the process begins around six months prior to the start of the season. Slot requests are often made in the form of series, which consist of multiple requests for the same time, and usually the same day of the week, over a period of at least five weeks. Airlines submit their slot requests using a standardized format that indicates for each request or request series, among other data, the requested times

for arrival at and departure from the airport, the dates for which the request applies, and any priority codes that may apply for the request.

The coordinator then proposes an initial allocation of slots to the airlines based on these requests. An airline may accept or reject an offer of a slot. After this initial allocation, the airlines and coordinators meet at the biannual IATA slot conference, whose purpose is to provide airlines with an opportunity to discuss schedule adjustments with coordinators and to facilitate bilateral trading of slots between airlines. Following the slot conference, a period of continuous slot allocation occurs until the start of the scheduling season, when airlines may make new requests or modify or delete existing requests.

The mechanism proposed in this paper aims to produce an initial allocation of slots following the rules that govern this. We now briefly describe these rules. For a more detailed summary of the IATA WSG for the initial allocation of slots, see Zografos, Salouras, and Madas (2012) and Ribeiro et al. (2018).

Each coordinated airport has a declared capacity that constrains how many slots can be allocated due to the availability of resources (e.g., runways, apron stands, etc.). This declared capacity is usually expressed in the form of rolling capacity constraints. These limit the number of arrival or departure slots or total number of slots that can be allocated over a given time interval (e.g., 15 minutes or 1 hour). The term *slot pool* is often used to refer to the availability of slots. It is important to bear in mind that in using this term, we do not refer to a single well-defined set of arrival and departure slots, but rather the possible sets of slots that can be allocated subject to the capacity constraints.¹ The schedule must also satisfy *turnaround constraints*, which prescribe that, in order to prepare an aircraft for the next flight, a departure must be scheduled a sufficient amount of time after its arrival.

The capacity constraints mean that not all requests can be allocated their preferred slots. Although the aim of the slot allocation is to construct a schedule that matches as far as possible the requested times, slots must also be allocated primarily according to the following main priority groups:

1. *Historic requests.* An airline that had a series of slots in the preceding season (and did not misuse them) is entitled to the same series of slots in the next season. Airlines with historic rights are also entitled to request a change to the time of a slot.

2. *New entrant requests.* An airline is given the status of *new entrant* if it does not already have significant presence at the airport. Of the remaining slots in the slot pool after historical requests have been allocated, 50% should be allocated to new entrants, unless slot requests by new entrants make up less than 50% of the remaining slots.

3. *Other requests.* All remaining requests have the lowest priority.

Note that the mechanism in this paper assumes that there are enough slots in each day to allocate to the requests. Hence, the 50% rule described above for new entrant requests can be disregarded for the purposes of this paper.

1.2. Literature Review

1.2.1. Slot-Allocation Models. Because of the complexity and size of the slot-allocation problem, a number of optimization models have been developed to allocate slots. A common concept used in these models is that of the *displacement* of a request, which is the absolute difference in time between the requested and allocated slots of a request. In order to satisfy slot requests as well as possible, these models typically aim to minimize some objective based on displacement. The most common objective is the *schedule displacement*, which refers to the total displacement across all slots requests.

The allocation of slots to incorporate the main features of the IATA WSG (scheduling for a season and request priorities) was first explicitly modelled as an optimization problem in the paper by Zografos, Salouras, and Madas (2012), who formulated it as an integer linear program (ILP). The model allocates requests for series of slots over a scheduling season, includes rolling capacity and turnaround constraints, and minimizes the total schedule displacement. The different priority classes of the requests are taken into account by solving the model *hierarchically*; that is, slots are allocated to the highest priority class by solving the slot-allocation model for only those requests, the rolling capacities are updated, the model is solved for the next highest priority group, and so on. The model is solved via a row generation algorithm using request data from three coordinated airports, demonstrating improvements over the actual allocations.

The model of Zografos and Jiang (2016, 2019) extends that of Zografos, Salouras, and Madas (2012) to incorporate fairness. A biobjective formulation is proposed to study the trade-off between schedule displacement and fairness, and is solved using the ϵ -constraint method. The model is solved both hierarchically and nonhierarchically for a congested airport. In the nonhierarchical case, priority classes are disregarded, and all requests are allocated slots simultaneously. The analysis reveals that in both cases, the fairness of an allocation could be substantially improved at the cost of only a small increase in schedule displacement.

Zografos, Androutsopoulos, and Madas (2018) present two biobjective formulations. The first formulation minimizes the schedule displacement and the largest single displacement allocated to a request.

The second formulation introduces acceptable time windows in which requests should be allocated slots and seeks to minimize the total schedule displacement and the number of requests allocated slots outside of their acceptable time windows. Airlines tend to be averse to large displacements, and so these models aim to construct a schedule that is more acceptable. These models were again solved using the ϵ -constraint method.

Ribeiro et al. (2018) present a slot-allocation model that uses a different formulation and that models the IATA WSG rules for the initial allocation to a greater level of detail than had been done previously. This model has four objectives, which are minimized lexicographically: the number of rejected requests, the maximum schedule displacement, the total schedule displacement, and the number of displaced requests. It also treats separately the IATA priority groups of “historical” and “changes to historical” requests (International Air Transport Association 2017, section 8.3, pp. 34–35). Finally, priority classes are taken into account by Ribeiro et al. (2018) by solving the model lexicographically rather than hierarchically.

Other models have been proposed for slot scheduling that do not incorporate the main aspects of the IATA WSG or that have not been solved for a scheduling season, but that incorporate other problem features. Jacquillat and Odoni (2015) propose a model that incorporates tactical operations such as runway configuration changes in order to calculate the expected queue lengths that arise from a given schedule. This allows capacity levels to be adjusted in such a way that queue lengths do not exceed a given level. This model is extended by Jacquillat and Vaze (2018) to ensure that displacement is apportioned fairly among the competing airlines. All the slot-allocation models discussed thus far concern the allocation of slots at a single airport. However, another important issue in the allocation of slots is ensuring that flights have compatible departure and arrival slots at their origin and destination airports. This can be achieved by solving a network-wide slot-allocation model, that is, a model that simultaneously allocates slots across a network of airports (Castelli, Pellegrini, and Pesenti 2012; Corolli, Lulli, and Ntamo 2014; Pellegrini et al. 2017; Benlic 2018). The sizes and computational complexities of both of these types of models mean that they cannot be solved exactly for a whole scheduling season.

1.2.2. Fairness. Fairness has been widely studied in the context of systems where limited resources or benefits must be shared between a group of participants. Although general axiomatic approaches to fairness exist (Nash 1950, Kalai and Smorodinsky 1975), there is no universally agreed upon definition, and it is often formulated based on problem context. Nevertheless,

there is generally a trade-off between fairness and total system efficiency.²

In the context of air traffic flow management (ATFM), schemes have been proposed for incorporating fairness into traffic management initiatives (TMIs) which reschedule flights in response to reductions in capacity of airports and airspace. A consensus has emerged that the most equitable way of allocating new arrival slots is to use the ration-by-schedule (RBS) rule, whereby flights are rescheduled on a first-scheduled, first-served basis (Wambsganss 1996). As this approach can lead to excessive delays and, in the case of multiple reductions of capacity, infeasible schedules (Barnhart et al. 2012), optimization models have been proposed that trade off overall delay against deviations from the RBS principle. Vossen et al. (2003) propose an optimization model to reschedule flights during a ground-delay program³ (GDP), which is found to reduce delays significantly while reducing variation in the average delays across the airlines. The paper by Barnhart et al. (2012) generalizes the RBS principle to account for the case of multiple reductions of capacity across a network, and proposes an optimization model for rescheduling flights that improves fairness according to this new definition. Bertsimas and Gupta (2015) present a model that reschedules flights across a whole airspace network. As well as minimizing overall delay, this model incorporates fairness by also minimizing reversals in the ordering of pairs of flights in RBS schedules.

Airport slot allocation differs from the ATFM interventions described above in that there is no obvious fair reference schedule. Instead, fairness has been defined by how schedule displacement is distributed across airlines requesting slots. Zografos and Jiang (2017) construct indices that measure how fairly each airline is treated in a given schedule. This fairness index is predicated on the principle that the proportion of total schedule displacement assigned to an airline should be similar to the proportion of requests made by it. Fairness is then incorporated into the slot-allocation model by adding constraints involving these fairness indices. The approach of Jacquillat and Vaze (2018), whose model is formulated and solved for a single day, is to incorporate fairness by lexicographically minimizing the schedule displacement normalized by the number of requests for each airline. Note that this is the axiomatic approach to fairness of Kalai and Smorodinsky (1975). A limitation of both of these approaches is that they treat all requests equally, regardless of when they are made. In particular, requests made for off-peak periods may cause airlines to be allocated more schedule displacement, despite the fact that these requests do not contribute to scheduling infeasibility.

1.2.3. Prioritization. The current IATA slot-allocation process and existing models do not consider explicitly the airline preferences regarding which requests should be displaced and by how much. By taking these into account, one could construct a schedule that is more acceptable to the participating airlines.

In ATFM, much work has been done on the involvement of airlines in the decision-making process to improve overall outcomes. One of the most famous examples of this, and which has been in operation for many years, is the schedule compression algorithm used in the GDP in the United States. This algorithm encourages airlines to submit accurate information regarding delays by allowing airlines preferential access to new slots when they declare flight delays or cancellations. The system has been demonstrated to reduce delays substantially compared with previous practice (Wambsganss 1996). Models have since been proposed that allow slots to be exchanged between airlines in more flexible ways and in more general settings (Vossen and Ball 2006, Bertsimas and Gupta 2015, Pilon et al. 2016). The works cited here use conditional trades, whereby an airline offers to give up one or more slots in exchange for gaining one or more other slots. There is trade-off with regard to the type of offers airlines are allowed to propose. The more general the type of offer, the more offers an airline can make, but the more computationally difficult it becomes to optimize the slot exchange. The less flexible the type of allowed offer, the fewer requests an airline can make, and the less effective the exchange.

Mechanisms have also been proposed that allow airlines to have a say on parameters of TMIs such as the GDP (e.g., scope and duration), by submitting their preferences on system-wide performance objectives (e.g., predictability and capacity). See, for instance, Swaroop and Ball (2013) and Evans, Vaze, and Barnhart (2016). All these schemes relate to TMIs that are implemented in the case of capacity reduction. A prioritization scheme has also been developed for the routing of flights across an airspace network (Dal Sasso et al. 2019). Here, an optimization model and heuristic algorithm are proposed, which attempt to minimize various objectives based on the deviation of a flight from the airline's ideal route and according to a ranking by each airline of all of its flights.

In the context of slot allocation, the incorporation of airline preferences is comparatively less well studied. A slot-exchange mechanism was proposed by Pellegrini, Castelli, and Pesenti (2012) to facilitate the secondary trading of slots after the initial allocation according to the IATA WSG (see Section 1.1). However, the model was formulated for only a single day, and relied on monetary transactions, which is inappropriate for the initial allocation of slots.

The idea of encoding preferences by weighting schedule displacement in slot-allocation models has also been explored. In this vein, Jacquillat and Vaze (2018) suggest that airlines could use weights to rank the importance of requests, or have an allocation of credit to apportion across all of their flights. The use of weights in the calculation of schedule displacement gives rise to two issues. First, in order to assign appropriate weights to requests, an airline should have some knowledge of which requests are likely to be displaced because of high demand. It would be wasteful for an airline to heavily weight requests that are not liable to be significantly displaced. Second, weighting requests does not give any guarantees to an airline regarding how much a request is displaced.

1.3. Contributions

The work in this paper builds on work presented previously in the extended abstract by Fairbrother and Zografos (2018). The main contribution of this paper is a two-stage mechanism for slot allocation that incorporates efficiency, fairness, and airline preferences, and that implements the main features of the IATA WSG (scheduling for a season and request priorities). In the first stage of the mechanism, a reference schedule is constructed that is fair with regard to how much displacement is allocated to each airline. In the second stage, airlines suggest adjustments to this reference schedule, and the schedule is adjusted to satisfy these requested changes as far as possible. The development of each of these stages constitutes a contribution in itself.

With regard to the first stage, a new fairness metric is proposed, which takes into account each airline's contribution to scheduling infeasibility. We call this *demand-based fairness*. In this way, the new fairness approach overcomes the problem with respect to the previous ones discussed above by not penalizing requests made for off-peak periods. We also identify the issue of *airline Pareto optimality* for (demand-based and non-demand-based) fairness metrics of this type and how the issue can be overcome. The definitions of peak and off-peak periods that are proposed for the slot-allocation problem also constitute a contribution.

The second stage of the mechanism is essentially a slot exchange. It is novel in that it is, to our knowledge, the first slot-exchange mechanism proposed for an entire scheduling season. It is also novel in how adjustments to the schedule are proposed by airlines. Rather than propose concrete conditional exchanges, airlines instead use a *displacement budget* in the form of a "credit" allocated to each airline to be used to describe how much displacement it would prefer each flight to have. The mechanism then adjusts the schedule to meet requested adjustments as far as possible subject to constraints that ensure airlines do not lose out by participating

in the mechanism. Such an approach allows airlines to flexibly specify their preferences in a simple way.

2. Summary of Overall Mechanism

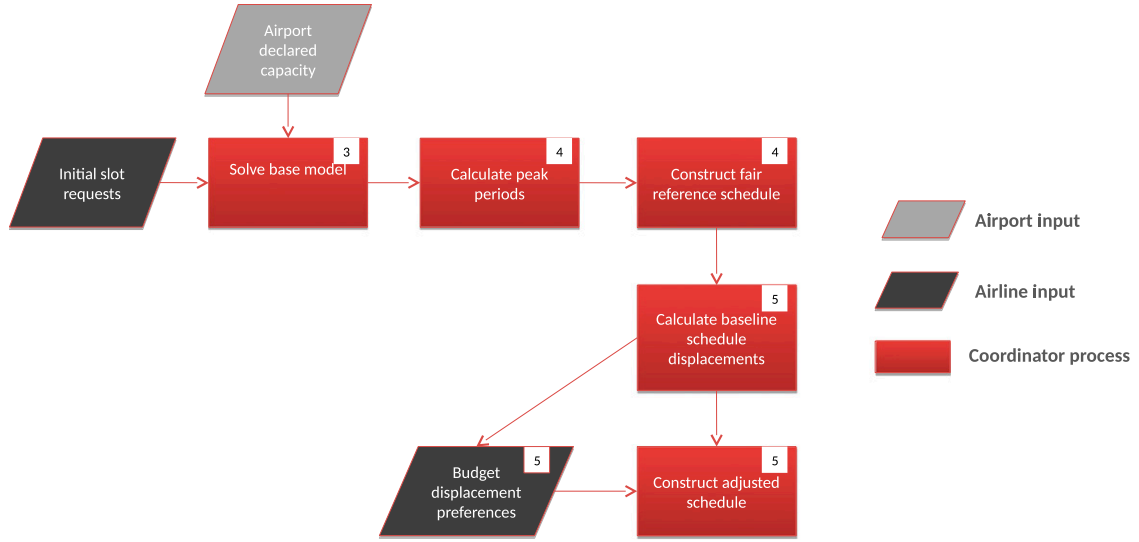
The mechanism proposed in this paper consists of two main stages. In the first stage, the mechanism constructs a fair reference schedule based on demand for slots in peak periods. During this stage, calculations must first be performed to determine for each period in the scheduling season whether demand is peak or off-peak. In the second stage of the mechanism, the airlines specify adjustments to the reference schedule, and the mechanism then adjusts this to meet as many of these preferences as possible. Our scheme is similar to that of Bertsimas and Gupta (2015), who propose a two-stage mechanism for assigning routes to flights in an ATFM problem. Their mechanism also constructs a reference schedule before adjusting it to better satisfy airline priorities.

The main aspects of our scheme are illustrated in Figure 1. Each process within the scheme is represented by a box, and this is annotated with the number of the section in which the process is detailed.

For the purpose of constructing a fair reference schedule, we propose a new approach to fairness. The approach taken in this paper is an enhancement of the proportionality principle of Zografos and Jiang (2016). In our approach, every period in the scheduling season is marked as either peak or off-peak, depending on whether there is any slack capacity, and our new fairness principle requires that the proportion of total displacement allocated to an airline be approximately equal to the proportion of requests it makes during peak periods.

The prioritization stage of our slot-allocation mechanism adjusts the reference schedule according to the airlines' preferences. The reference schedule is used to calculate for each request a *baseline displacement*, and the total baseline displacement for each airline is referred to as its displacement budget. The airlines must specify for each request a *preferred displacement* such that the sum of the preferred displacements of their requests is at least equal to its displacement budget. In this way, airlines can specify a smaller preferred displacement than the baseline for requests that should be prioritized in return for specifying larger preferred displacements for requests that are not a priority. The mechanism then adjusts the reference schedule so that the requests are allocated displacements as close to the airlines' preferred values as possible. This is done in a way that guarantees that the displacement of a request will not be greater than the airline's preferred displacement, or the baseline displacement if this is larger.

Unlike the method of weighting the displacement, as suggested by Jacquillat and Vaze (2018) as a way of

Figure 1. (Color online) Overall Slot-Allocation Scheme

prioritizing requests, our prioritization mechanism is done with respect to a reference schedule, and therefore the airlines can specify their preferences in the knowledge of which of their requests are at risk for a large displacement. Our prioritization mechanism also does not require the airlines to reveal any of their operating costs to the slot coordinator. Finally, our mechanism allows airlines to specify their preferences in a much more flexible way than other schedule adjustments approaches, such as the at-most-at-least trade offers used by Bertsimas and Gupta (2015) where an airline specifies it would be willing to reschedule one flight to a later time in exchange for moving another flight to an earlier time.

3. Base Slot-Allocation Model

In this section, we define key terminology and notation related to our approach and present an integer linear program that forms the core of our mechanism and which is extended in later sections.

3.1. Basic Concepts and Terminology

The slot-allocation problem takes place over a scheduling season. The scheduling season consists of a set of days $d \in \mathcal{D}$, with each day divided into a set of equally sized (coordination) time intervals $t \in \mathcal{T} = [0, T - 1]$. For the airport that we use for numerical testing in Section 7, the coordination time interval has a length of 15 minutes, but for many other airports it may be 5 or 10 minutes. We reserve the term *time period* to mean a pair $(d, t) \in \mathcal{D} \times \mathcal{T}$, which corresponds to a time interval for a particular day in the scheduling season. An arrival (departure) slot is a permission to use airport infrastructure for the purposes of landing (take-off) during a given time period (d, t) .

For the purposes of this paper, a *request series*, $m \in \mathcal{M}$, is a set of requests made by an airline for a

series of arrival (or departure) slots for the time interval $t_m \in \mathcal{T}$, on days $\mathcal{D}_m \subseteq \mathcal{D}$, and by *request*, we refer to an individual request for an arrival (or departure) slot that is made as part of a series. All requests in a single request series must be allocated slots for the same time interval. The purpose of request series is thus to allow airlines to create consistent schedules. Note that some request series may consist of only a single request. An individual request is referred to by a pair of indices (m, d) , where $m \in \mathcal{M}$ and $d \in \mathcal{D}_m$. Note that our definition of request series is more general than that given in the IATA WSG, where a request series comprises of requests of a period over at least five weeks. Our definition is only necessary for the purposes of exposition, and we process requests according to the definition in the IATA WSG.

A *request pair*, $(m_1, m_2) \in \mathcal{P}$, consists of an arrival request series, m_1 , and a subsequent departure request series, m_2 , which are to be operated by the same aircraft and which must be allocated slots compatible with an aircraft's turnaround time. Specifically, for a request pair $(m_1, m_2) \in \mathcal{P}$, the absolute difference between the time intervals of the slots allocated to the arrival and departure must be greater than the *minimum turnaround time*, $l_{m_1 m_2}$, and less than the *maximum turnaround time*, $l_{m_1 m_2}$. We assume throughout that all request series are made as part of a request pair.

The number of slots that can be allocated over a sequence of consecutive time periods is limited by the *rolling capacity constraints*. For a duration c and time period (d, t) , the number of arrival slots that can be allocated for the time intervals $\{t, \dots, t + c - 1\}$ on day d is at most α_c^{dt} . Similarly, the number of departure slots and the total number of arrival and departure slots allocated in these periods are at most β_c^{dt} and γ_c^{dt} , respectively.

A *schedule* is an allocation of slots to the set of request series $\mathcal{M}$ and is *feasible* if it satisfies airport capacity and turnaround constraints. By displacement we mean the absolute difference measured between the requested time intervals of a request series and the slots allocated to that request series; that is, if a request series m is allocated slots at time interval t , then the displacement, denoted by f_m^t , is $|\mathcal{D}_m||t - t_m|$. The total displacement of a schedule, or schedule displacement, is the sum of displacements over all requests. The aim of the base slot-allocation model is to construct a feasible schedule that minimizes total displacement.

3.2. Optimization Model

We present a modified version of the model by Zografos, Salouras, and Madas (2012). The sets, parameters, and variables used in this model and its extensions are given in Table 1. The full optimization model is given in (1)–(9). We call this the *base model*.

The main decision variables in this model, x_m^t , are binary and indicate whether a request series $m \in \mathcal{M}$ is assigned slots at $t \in \mathcal{T}$.

$$\text{Minimize } \sum_{m \in \mathcal{M}} \sum_{t \in \mathcal{T}} f_m^t x_m^t \quad (1)$$

$$\text{subject to } \sum_{t \in \mathcal{T}} x_m^t = 1, \quad m \in \mathcal{M}, \quad (2)$$

$$y_{dt}^A = \sum_{m \in \mathcal{M}^A} a_m^d x_m^t, \quad t \in \mathcal{T}, \quad d \in \mathcal{D}, \quad (3)$$

$$y_{dt}^D = \sum_{m \in \mathcal{M}^D} a_m^d x_m^t, \quad t \in \mathcal{T}, \quad d \in \mathcal{D}, \quad (4)$$

$$\sum_{t \in [s, s+c-1]} y_{dt}^A \leq \alpha_c^{ds}, \quad c \in \mathcal{C}, \quad d \in \mathcal{D}, \quad s \in [0, T-c], \quad (5)$$

$$\sum_{t \in [s, s+c-1]} y_{dt}^D \leq \beta_c^{ds}, \quad c \in \mathcal{C}, \quad d \in \mathcal{D}, \quad s \in [0, T-c], \quad (6)$$

$$\sum_{t \in [s, s+c-1]} y_{dt}^A + y_{dt}^D \leq \gamma_c^{ds}, \quad c \in \mathcal{C}, \quad d \in \mathcal{D}, \quad s \in [0, T-c], \quad (7)$$

$$l_{m_1, m_2} \leq \sum_{t \in \mathcal{T}} t x_{m_1}^t - \sum_{t \in \mathcal{T}} t x_{m_2}^t \leq \bar{l}_{(m_1, m_2)}, \quad (m_1, m_2) \in \mathcal{P}, \quad (8)$$

$$x_m^t \in \{0, 1\}, \quad m \in \mathcal{M}, \quad t \in \mathcal{T}. \quad (9)$$

The objective in (1) minimizes the total displacement of the constructed schedule. The constraints (2) ensure that each request series is allocated a set of slots. The constraints (3) and (4) define the auxiliary variables y_{dt}^A and y_{dt}^D to be, respectively, the number of arrival and departure slots allocated for time period (d, t) . Although this model could be formulated without these, we include them for notational convenience. These are used in Section 4 for the characterization of peak and off-peak periods.

The constraints (5)–(7) enforce rolling capacity constraints. In particular, constraints (5) limit the rolling

Table 1. Notation Used for the Single Airport Slot-Allocation Model

Symbol	Definition
Sets	
$\mathcal{A}$	Set of airlines
$\mathcal{M}$	Set of request series
$\mathcal{M}^{A(D)} \subset \mathcal{M}$	Set of all arrival (departure) requests series
$\mathcal{M}_a \subset \mathcal{M}$	Set of requests series made by airline $a \in \mathcal{A}$
$\mathcal{M}_a^{A(D)} \subset \mathcal{M}^{A(D)}$	Set of arrival (departure) request series made by airline a
$\mathcal{D}$	Set of days in scheduling season
$\mathcal{D}_m \subseteq \mathcal{D}$	Set of days for which request series m must be scheduled
$\mathcal{P} \subset \mathcal{M} \times \mathcal{M}$	Set of request series pairs (m_a, m_d) where m_d corresponds to the departure of an aircraft following the arrival m_a
$\mathcal{C}$	Set of durations of rolling capacity constraints
$\mathcal{T} = \{1, \dots, T\}$	Set of coordination time intervals
Parameters	
$t_m \in \mathcal{T}$	Requested time for request series m
r_a	Proportion of peak requests made by airline a
f_m^t	Displacement from assigning slot t to request series m
l_p ($\bar{l}_p$)	Minimum (maximum) turnaround time for request pair $p \in \mathcal{P}$
a_m^d	Binary value indicating whether request series $m \in \mathcal{M}$ includes a request for day $d \in \mathcal{D}$
α_c^{ds}	Rolling capacity on arrivals for period $[s, s+c-1]$ on day $d \in \mathcal{D}$
β_c^{ds}	Rolling capacity on departures for period $[s, s+c-1]$ on day $d \in \mathcal{D}$
γ_c^{ds}	Rolling capacity on total movements for period $[s, s+c-1]$ on day $d \in \mathcal{D}$
Decision variables	
x_m^t	Indicates whether request series m is assigned slots at t
y_{dt}^A	Aggregate number of arrival slots allocated for period (d, t)
y_{dt}^D	Aggregate number of departure slots allocated for period (d, t)

number of arrivals, constraints (6) limit the rolling number of departures, and constraints (7) limit the rolling total number of movements. For notational simplicity, we assume that all three types of capacity constraints use the same set of durations. This leads to no loss of generality, as we can simply take the capacity to be infinity in cases where there is no limit specified for a particular duration and movement type. Finally, constraints (8) bound the turnaround time for each request pair $(m_1, m_2) \in \mathcal{P}$.

4. Construction of a Fair Reference Schedule

In this section, we present a new approach to fairness that distinguishes between peak and off-peak time periods and explain how this is used to construct a reference schedule for our slot-allocation mechanism. In Sections 4.1 and 4.2, we define the related concepts of slackness, saturation, and peak and off-peak periods. In Section 4.3, we use peaks periods to define a new fairness metric. In Section 4.4, we describe a technical issue related to incorporating fairness into the base model. Finally, in Section 4.5, we present an algorithm for constructing a reference schedule.

4.1. Slackness and Saturation

Previously, we said that a peak period could be thought of as one where there is no slack capacity, that is, where no extra slots for the period can be allocated without breaking a capacity constraint. We call such a period *saturated*. We now give precise mathematical definitions to these concepts. Note that because arrival and departure slots are subject to different capacity constraints, slackness and saturation must be defined differently for each type of runway movement.

Definition 1. Suppose $\tilde{\mathbf{y}} := (\tilde{y}_{dt}^M)_{d \in \mathcal{D}, t \in \mathcal{T}}$ is the aggregate schedule corresponding to a (possibly infeasible) solution to problem (1)–(9). The arrival, departure, and total movement slacks at a time period (d, t) are defined to be, respectively,

$$\begin{aligned}\bar{\alpha}_{dt}(\tilde{\mathbf{y}}) &= \min_{c \in \mathcal{C}} \min_{s \in [\max(t-c+1, 0), t]} \left(\alpha_c^{dt} - \sum_{\bar{s}=s}^{s+c-1} \tilde{y}_{d\bar{s}}^A \right)_+, \\ \bar{\beta}_{dt}(\tilde{\mathbf{y}}) &= \min_{c \in \mathcal{C}} \min_{s \in [\max(t-c+1, 0), t]} \left(\beta_c^{dt} - \sum_{\bar{s}=s}^{s+c-1} \tilde{y}_{d\bar{s}}^D \right)_+, \\ \bar{\gamma}_{dt}(\tilde{\mathbf{y}}) &= \min_{c \in \mathcal{C}} \min_{s \in [\max(t-c+1, 0), t]} \left(\gamma_c^{dt} - \sum_{\bar{s}=s}^{s+c-1} (\tilde{y}_{d\bar{s}}^A + \tilde{y}_{d\bar{s}}^D) \right)_+, \end{aligned}$$

where $u_+ = \max(u, 0)$.

The arrival slack gives the number of extra arrival slots for time period (d, t) that can be allocated before an arrival constraint is broken. Departure and total movement slack can be interpreted similarly.

Definition 2. Suppose $\tilde{\mathbf{y}}$ is an aggregate schedule. Then, the corresponding *saturation* indices are defined as follows:

$$\begin{aligned}A_d^t(\tilde{\mathbf{y}}) &= \begin{cases} 1 & \text{if } \bar{\alpha}_{dt}(\tilde{\mathbf{y}}) = 0 \text{ or } \bar{\gamma}_{dt}(\tilde{\mathbf{y}}) = 0, \\ 0 & \text{otherwise;} \end{cases} \\ D_d^t(\tilde{\mathbf{y}}) &= \begin{cases} 1 & \text{if } \bar{\beta}_{dt}(\tilde{\mathbf{y}}) = 0 \text{ or } \bar{\gamma}_{dt}(\tilde{\mathbf{y}}) = 0, \\ 0 & \text{otherwise.} \end{cases} \end{aligned}$$

A period (d, t) is said to be *arrival saturated* with respect to $\tilde{\mathbf{y}}$ if $A_d^t(\tilde{\mathbf{y}}) = 1$. Similarly (d, t) is *departure saturated* with respect to $\tilde{\mathbf{y}}$ if $D_d^t(\tilde{\mathbf{y}}) = 1$.

A period is arrival (departure) saturated if no more arrival (departure) slots can be allocated without breaking a capacity constraint.

4.2. Characterization of Peak Periods

A natural definition for arrival (departure) peak periods is those that are arrival (departure) saturated with respect to the *default schedule*. By default schedule, we mean the schedule where every request series is allocated its preferred slots. Mathematically, the default schedule is characterized by the following solution:

$$\tilde{x}_m^t = 1 \Leftrightarrow t = t_m \text{ for all } m \in \mathcal{M}.$$

The periods that are saturated with respect to the default schedule are those periods where the number of requests equals or exceeds the available capacity. However, this definition for peak periods is flawed. The default schedule will typically be infeasible. (If it were, feasible there would be no need to solve the slot-allocation problem.) When constructing a feasible schedule, some requests will be displaced, which may lead to other periods becoming saturated. These periods cannot accommodate any extra slots without breaking capacity constraints, and so should be considered to be peak. A definition for peak periods that overcomes this issue is the following.

Definition 3 (Period Sensitivity). A period (d, t) is said to be arrival (departure) sensitive if an single extra arrival (departure) request for that period increases the optimal total displacement as calculated in problem (1)–(9).

Given that the problem (1)–(9) is an integer program with a complex structure, the calculation of whether each period is arrival or departure sensitive would largely have to be done by brute force; that is, for each day and time period, we would have to adjust the requests and resolve the problem. We will instead use an approximation of sensitive periods to define peak and off-peak periods, which can be calculated easily after solving the problem (1)–(9) once, and which is conservative, in the sense that any period that is sensitive will be deemed a peak period.

Definition 4. Let $(\mathbf{x}^*, \mathbf{y}^*)$ be an optimal solution to the slot-allocation problem (1)–(9). Then, a period (d, t) is an arrival peak period if it is arrival saturated with respect to $\mathbf{y}^*$. Similarly, (d, t) is a departure peak period if it is departure saturated with respect to $\mathbf{y}^*$.

The downside of this definition of peak periods is that it depends on the optimal solution, which may not be unique. However, as well as having the advantage of being easy to calculate, the next result shows that for any optimal solution, the arrival (departure) peak periods must cover all arrival-sensitive (departure-sensitive) periods.

Proposition 1. Suppose the time period (d, t) is arrival (departure) sensitive; then $A_d^t(\mathbf{y}^*) = 1$ ($D_d^t(\mathbf{y}^*) = 1$) for any optimal solution $(\mathbf{x}^*, \mathbf{y}^*)$ of the base problem (1)–(9).

The proof of this result can be found in Online Appendix A. Unfortunately, the converse of the result is not true, and so we also give a counterexample to this in the same appendix.

One could imagine a peak period indicator that indicated not only the presence of peak demand but also the severity. A possible way of measuring this would be to look at the actual change in optimal solution value rather than simply whether it increases or not. However, such an approach would likely be very computationally expensive, and we are not aware of any numerically efficient approximation approaches.

4.3. Demand-Based Fairness

For notational convenience, unless it is necessary, from now on we will suppress the dependency on an aggregate schedule $\mathbf{y}^*$ from our notation for peak periods and simply write them as A_d^t and D_d^t . Using the decision variables and parameters from the base model, the amount of schedule displacement for each airline can be expressed as follows:

$$s_a = \sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} f_m^t x_m^t = \sum_{t \in \mathcal{T}} |\mathcal{D}_m| |t - t_m| x_m^t, \text{ with} \\ S = \sum_{a \in \mathcal{A}} s_a$$

for the aggregate displacement across airlines, in other words, the total displacement.

Given arrival and departure peak indicators, the number of peak requests made by an airline a is given by

$$C_a := \sum_{m \in \mathcal{M}_a^A} \sum_{d \in \mathcal{D}_m} A_d^{t_m} + \sum_{m \in \mathcal{M}_a^D} \sum_{d \in \mathcal{D}_m} D_d^{t_m},$$

and the proportion of peak requests made by airline a is

$$r_a := \frac{C_a}{\sum_{a \in \mathcal{A}} C_a},$$

assuming that $\sum_{a \in \mathcal{A}} C_a \neq 0$.

We propose that the proportion of total displacement allocated to an airline should be similar to its proportion of peak requests. In this way, airlines will not be penalized for requests made during off-peak periods. We call this the *peak request proportionality principle*. The new fairness index for an airline is defined to be the proportion of total displacement allocated to an airline divided by its proportion of peak requests, with special cases to account for the cases where an airline makes no peak requests:

$$\mu_a := \begin{cases} \frac{s_a}{r_a} & \text{if } r_a \neq 0, \\ 1 & \text{if } r_a = 0 \text{ and } s_a = 0, \\ \infty & \text{if } r_a = 0 \text{ and } s_a \neq 0. \end{cases} \quad (10)$$

Based on the peak request proportionality principle, μ_a should be interpreted as follows:

$$\begin{aligned} \mu_a = 1.0, & \quad \text{airline } a \text{ is fairly treated,} \\ \mu_a < 1.0, & \quad \text{airline } a \text{ is favoured,} \\ \mu_a > 1.0, & \quad \text{airline } a \text{ is disfavoured.} \end{aligned}$$

We refer to these indices as demand-based fairness indices. Note that it is possible to construct other fairness indices based on the peak request proportionality principle. We use this one as it is similar to indices already used in the literature, and the value of the index has the following easy interpretation: if $\mu_a = q$, then airline a receives q times more (or less) total schedule displacement than it should. We do not claim that this is necessarily the best demand-based fairness index that could be used, but only aim to demonstrate the advantages of demand-based fairness over non-demand-based fairness.

Fairness metrics can be constructed from fairness indices to measure the overall fairness of a schedule. Throughout this paper, we use a fairness metric proposed by Zografos and Jiang (2017) called the maximum deviation from absolute fairness (MDA). This is defined as follows:

$$\mu_{MDA} = \max_{a \in \mathcal{A}} |\mu_a - 1|. \quad (11)$$

This metric measures fairness by focusing on the worst case. Again, it is possible to construct other fairness metrics from the airline fairness indices, but this is not the focus of this paper. Two other fairness metrics constructed from fairness indices were proposed by Zografos and Jiang (2017), namely, the deviation from average fairness and the Gini coefficient.

Solving the base model will not necessarily yield a schedule that is fair according to the peak request proportionality principle. Fairness can be incorporated into the slot-allocation model through the minimization of a fairness metric. In the case of the MDA

metric, this is a nonlinear function with respect to the decision variables (x_m^t), which could render the optimization model intractable if they were included directly in the objective function. This problem can be circumvented by using these metrics in constraints, as these can be linearized as follows:

$$\mu_{MDA} \leq \epsilon \iff |s_a - Sr_a| \leq \epsilon r_a S, \quad a \in \mathcal{A}. \quad (12)$$

We call the base model with the above fairness constraints the ϵ -fairness model.

In the case where $r_a = 0$, the constraint above becomes $s_a = 0$. The next result shows that constraint (12) is satisfiable in this case, as long as the airline requests arrival and departure slots which are compatible with the required turnaround time.

Proposition 2. Suppose $(\mathbf{x}^*, \mathbf{y}^*)$ is an optimal solution to problem (1)–(9). Let (A_d^t) and (D_d^t) be the arrival and departure saturation indicators for $\mathbf{y}^*$, and define $\tilde{\mathcal{M}} = \bigcup_{a \in \mathcal{A}: r_a = 0} \mathcal{M}_a$. Assume the following conditions hold:

- For each $(m_1, m_2) \in \mathcal{P}$, we have $L_{m_1 m_2} \leq t_{m_2} - t_{m_1} \leq \bar{L}_{m_1 m_2}$.
- For each $(m_1, m_2) \in \mathcal{P}$, if $m_1 \in \tilde{\mathcal{M}}$, then $D_d^{t_{m_2}} = 0$ for all $d \in \mathcal{D}_{m_2}$.
- For each $(m_1, m_2) \in \mathcal{P}$, if $m_2 \in \tilde{\mathcal{M}}$, then $A_d^{t_{m_1}} = 0$ for all $d \in \mathcal{A}_{m_1}$.

Then $\mathbf{x}^*$ satisfies $s_a = 0$ for all $a \in \mathcal{A}$ such that $r_a = 0$.

The proof for this proposition can again be found in Online Appendix A. Assumption (i) of Proposition 2 holds as long as all airlines make requests consistent with the required turnaround times for the aircraft.

Assumptions (ii) and (iii) automatically hold if the request series of every request pair is made by the same airline, that is, if $(m_1, m_2) \in \mathcal{P}$ implies that $m_1, m_2 \in \mathcal{M}_a$ for some $a \in \mathcal{A}$. In this case, for $(m_1, m_2) \in \mathcal{P}$, we have $m_1 \in \tilde{\mathcal{M}}$ if and only if $m_2 \in \tilde{\mathcal{M}}$.

Because we compare demand-based fairness with the previously proposed non-demand-based fairness approach of Zografos and Jiang (2019) in our numerical experiments, we now define the latter. For an airline a , the non-demand-based fairness index ρ_a is defined to be the proportion of schedule displacement allocated to that airline divided by the proportion of requests made by the airline:

$$\rho_a = \frac{\frac{S_a}{S}}{\frac{\sum_{m \in \mathcal{M}_a} |\mathcal{D}_m|}{\sum_{m \in \mathcal{M}} |\mathcal{D}_m|}}. \quad (13)$$

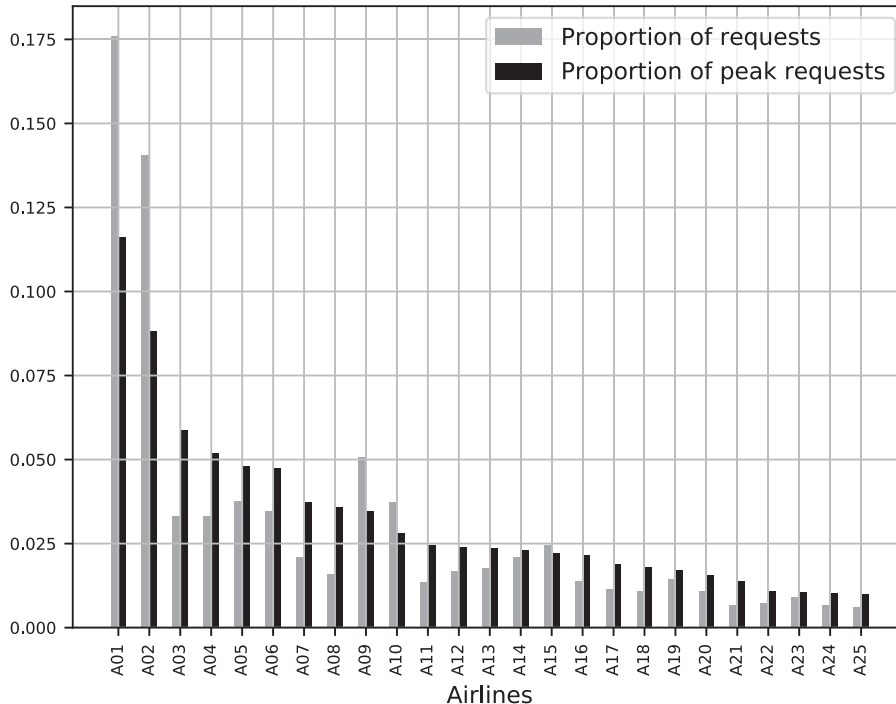
This can be interpreted in the same way as μ_a . Similarly, the non-demand-based MDA metric is defined to be

$$\rho_{MDA} = \max_{a \in \mathcal{A}} |\rho_a - 1|.$$

This can be linearized when used in a constraint in the same way as the demand-based MDA metric.

To demonstrate the importance of using demand-based fairness over non-demand-based fairness, we plot the proportions of requests and peak requests for each airline for some real request data (to be described in more detail in Section 7.1) in Figure 2. Because of the very large number of airlines operating at the airport (80), for clarity, we plot these proportions only for airlines who have a proportion of peak requests

Figure 2. Proportions of Requests and Peak Requests for Major Airlines at a Coordinated Airport



greater than 0.01. This plot shows that for many airlines, the proportions of requests and peak requests can differ greatly, meaning that some airlines would be significantly penalized for their off-peak requests under a non-demand-based fairness approach.

4.4. Airline Pareto Optimality

Solving the base slot-allocation problem with fairness constraints (12) yields a solution that is (weakly) Pareto optimal with respect to objectives of total displacement (Equation (1)) and MDA (Equation (11)). In this section, we consider a different type of Pareto optimality.

Suppose $\tilde{x}$ and $\bar{x}$ are two feasible schedules. Then, we say that the schedule $\tilde{x}$ dominates $\bar{x}$ with respect to airline displacements if

$$\sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} f_m^t \tilde{x}_m^t \leq \sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} f_m^t \bar{x}_m^t \quad \text{for all } a \in \mathcal{A},$$

with at least one inequality strict; that is, one schedule dominates another if the displacements allocated to each airline are less than or equal to those of the other schedule, and where the displacement is strictly less for at least one airline. We say that a feasible schedule is *airline Pareto optimal* if it is not dominated with respect to airline displacements by any other feasible schedule. If a schedule is airline Pareto optimal, no airline's displacement can be reduced without increasing it for another.

It is a natural requirement for a fairness scheme to yield a solution that is Pareto optimal in this sense (Nash 1950, Kalai and Smorodinsky 1975), that is, with respect to the objectives of the participants in the system. A schedule that is not airline Pareto optimal allocates more displacement than necessary to some airline, and so will likely be unacceptable. Unfortunately, the use of both demand and nondemand fairness constraints with the base model will not necessarily yield an airline-Pareto-optimal schedule. The following toy example demonstrates that fairness constraints that are too tight can yield a schedule that is not airline Pareto optimal.

Example 1. Suppose that we have two airlines each requesting a slot for the same time period, for a single day, and where there is a total slot limit of one for each

time period. The optimal solution in this case is to displace one of the airline's requests by one time period, which results in an MDA value (for both demand- and non-demand-based fairness indices) of 1. If we include a fairness constraint where we require the MDA value to be less than one, then the only way to satisfy this is to displace both requests. This is illustrated in Figure 3. Clearly, this is not an airline-Pareto-optimal solution.

Fortunately it is easy to check whether a feasible schedule is airline Pareto optimal. The process for doing this is presented in Algorithm 1.

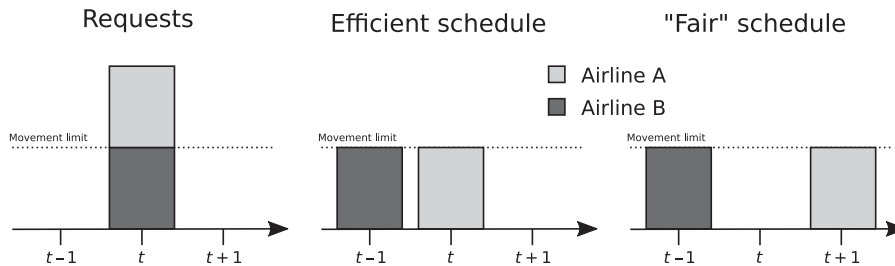
Algorithm 1 (Checks Whether a Given Feasible Schedule Is Airline Pareto Optimal)

Input: $\bar{x}$ feasible schedule

1. Construct base slot-allocation model;
2. **for** $a \in \mathcal{A}$ **do**
3. $\Sigma_a \leftarrow \sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} f_m^t \bar{x}_m^t$;
4. Add constraint $\sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} f_m^t x_m^t \leq \Sigma_a$;
5. Solve slot-allocation model and let $\tilde{x}$ be the optimal solution;
6. **for** $a \in \mathcal{A}$ **do**
7. **if** $\sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} f_m^t \tilde{x}_m^t < \Sigma_a$, **then**
8. **return** False;
9. **return** True;

The algorithm first calculates the schedule displacement for each airline and then constructs a new base slot-allocation problem where the displacement for each airline is bounded by these values. Optimizing this new model yields a new feasible schedule whose airline displacements do not exceed the previous airline displacements. The algorithm then checks for each airline whether its displacement is less than the previous displacement. If this is the case, then the new solution dominates the previous one, so the original schedule is not airline Pareto optimal and the algorithm returns False. If the new displacements for all airlines are the same as the previous displacements, then no feasible schedule that dominates the original exists, so the original schedule is airline Pareto optimal, and the algorithm returns True.

Figure 3. Toy Example to Illustrate How Schedules That Are Not Airline Pareto Optimal Arise from Use of Fairness Constraints



4.5. Efficient Frontier and Selection of a Reference Schedule

There is a trade-off between the total displacement and the fairness of a schedule. To examine this trade-off, we can calculate the efficient frontier between these two conflicting objectives using the ϵ -constraint method. After calculating the efficient frontier between total displacement and fairness, one has potentially several schedules from which to choose a reference schedule. Following our discussion in Section 4.4, we first eliminate from consideration all schedules that are not airline Pareto optimal. As explained above, a non-airline-Pareto-optimal schedule will likely be unacceptable to the participating airlines. In addition, an airline-Pareto-optimal solution is required to use the displacement budget mechanism described in Section 5.

To select a schedule from the available airline-Pareto-optimal schedules, we use a simpler scheme, which we refer to as the *acceptable price of fairness*. Denote by S^* the optimal value to the base model, and by $S^*(\epsilon)$ the optimal value for the ϵ -fairness problem. Then the price of fairness of the optimal solution yielded by the latter problem is defined to be

$$\frac{S^*(\epsilon) - S^*}{S^*};$$

that is, it is the relative deterioration in the optimal total displacement when introducing an MDA constraint with right-hand side ϵ . For $\alpha > 0$, the α -acceptable price of fairness criterion selects the most fair (airline-Pareto-optimal) schedule that has a price of fairness of less than α . The parameter α is chosen according to the airport coordinator's and airlines' tolerances to a higher total displacement. This criterion for selecting a fair solution was suggested by Bertsimas, Farias, and Trichakis (2012), who apply it to a different fairness scheme that also depends on the specification of a fairness parameter.

Algorithm 2 describes the overall procedure for calculating the fair reference schedule given the airline requests and airport capacities. It first calculates peak periods, before running a modified version of the ϵ -constraint method, which terminates once the price of fairness has been exceeded. Note that at each iteration of this algorithm ϵ is decremented by a factor δ rather than by a fixed value. We choose to use this rule, as the optimal total displacement and required solution time become a lot more sensitive to changes in the value of ϵ when ϵ is small.

Algorithm 2 (Calculation of Reference Schedule)

Input: Airline requests, airport coordination parameters, $0 < \delta < 1$ decrement rate of ϵ , α acceptable price of fairness

Output: $\bar{x}$ reference schedule

1. Solve base model;
2. $(\bar{x}, \bar{y}) \leftarrow$ optimal solution, $S \leftarrow$ optimal total displacement;
3. Calculate peak periods (A_d^t) and (D_d^t) from $\bar{y}$;
4. Calculate fairness indices and initial MDA ϵ using (10) and (11);
5. **repeat**
6. $\epsilon \leftarrow \delta\epsilon$;
7. Solve ϵ -constraint problem;
8. **if** Problem infeasible, **then**
9. | Break;
10. **else if** Optimal solution found, **then**
11. Set $\tilde{x}, S_\epsilon$ be optimal schedule and optimal solution value;
12. **if** $\frac{S_\epsilon - S}{S} > \alpha$, **then**
13. | break;
14. **if** $\tilde{x}$ is airline Pareto optimal, **then** $\bar{x} \leftarrow \tilde{x}$;
15. **return** $\bar{x}$;

Although other rules for selecting a reference schedule are possible, the acceptable price of fairness rule has two major advantages. First, unlike a voting process, the procedure does not require any extra input from the airlines. Second, the method does not require one to calculate the entire efficient frontier, which can save considerable computation time. This is because the ϵ -constraint method used to calculate the efficient frontier consists of solving the ϵ -fairness problem for a decreasing sequence of values of ϵ until the problem becomes infeasible. As the fairness constraints become tighter, the optimal total displacement increases. Therefore, once the price of fairness has been exceeded, it will be exceeded for all smaller values of ϵ , and so there is no need to solve more fairness problems. As will be seen in the numerical tests in Section 7, the time saved from not having to solve the ϵ -fairness problem for small values of ϵ is particularly large.

5. The Displacement Budget Mechanism

The incorporation of fairness into the slot-allocation model ensures that each airline is allocated a reasonable amount of displacement. However, there may be some flexibility in how the displacement assigned to an airline may be distributed among its requests. The acceptability of a schedule can therefore be improved by taking into account an airline's preferences regarding which requests should be displaced and by how much. We propose a mechanism, which we call the *displacement budget mechanism*, that adjusts the reference schedule in order to better satisfy the airlines' displacement preferences. Throughout this section, we use $\bar{x}$ to denote the reference schedule that can be calculated using Algorithm 2.

In Section 5.1, we describe the overall displacement budget mechanism. In Section 5.2, we describe

constraints on how airlines specify their preferences. Finally, in Section 5.3, we present a model to facilitate the specification of airline preferences for use in the displacement budget mechanism.

5.1. Overall Model

The fair reference schedule is used to determine the baseline displacement for each request in the reference schedule, namely,

$$\sigma_m := \sum_{t \in \mathcal{T}} f_m^t \bar{x}_m^t.$$

The displacement budget for an airline $a \in \mathcal{A}$ is defined to be the total baseline displacement for all of the airline's requests, that is,

$$\Sigma_a := \sum_{m \in \mathcal{M}_a} \sigma_m.$$

The mechanism requires airlines to specify, for each request m , a preferred displacement δ_m , namely, an expressed preference that the request be scheduled in the set

$$\{t \in \mathcal{T} : f_m^t \leq \delta_m\}.$$

The preferences must be specified so as to satisfy the *displacement budget constraint*:

$$\sum_{m \in \mathcal{M}_a} \delta_m \geq \Sigma_a. \quad (14)$$

The displacement budget thus represents the amount of displacement an airline should be allocated, and the mechanism requires airlines to specify their preferences for how this should be allocated among its requests. If an airline specifies a preferred displacement less than the baseline displacement for some requests, it must specify greater preferred displacements for others. Besides the displacement budget constraint, other constraints apply to the setting of preferred displacements, and these are discussed in Section 5.2. A systematic method that airlines could use for setting preferred costs is presented in Section 5.3.

A request for which an airline specifies a smaller preferred displacement than the baseline displacement is called an *improvement request*. We refer to all other requests as *nonimprovement requests*. The aim of the displacement budget mechanism is to adjust the reference schedule to satisfy improvement requests as far as possible.

The mechanism works by solving a modified version of the base model. For each $m \in \mathcal{M}$, we introduce variables π_m to measure the change in displacement with respect to the reference schedule. For improvement requests, we also introduce variables w_m ,

which measure the amount of displacement by which preferred displacements are missed:

$$\text{lexmin} \left(\sum_{m \in \mathcal{M} : \delta_m < \sigma_m} w_m, \sum_{m \in \mathcal{M}} \sum_{t \in \mathcal{T}} f_m^t x_m^t \right) \quad (15)$$

subject to constraints (2)–(9)

$$\pi_m = \sum_m f_m^t x_m^t - \sigma_m \text{ for all } m \in \mathcal{M}, \quad (16)$$

$$w_m \geq (\sigma_m - \delta_m) + \pi_m \text{ for each } m \in \mathcal{M} \\ \text{such that } \delta_m < \sigma_m, \quad (17)$$

$$\sum_{m \in \mathcal{M}_a} \pi_m \leq \nu \Sigma_a \text{ for each } a \in \mathcal{A}, \quad (18)$$

$$\pi_m \leq 0 \text{ for each } m \in \mathcal{M} \text{ such that } \delta_m < \sigma_m, \quad (19)$$

$$\pi_m \leq \delta_m - \sigma_m \text{ for each } m \in \mathcal{M} \\ \text{such that } \delta_m \geq \sigma_m, \quad (20)$$

$$w_m \geq 0 \text{ for each } m \in \mathcal{M} \text{ such that } \\ \delta_m < \sigma_m. \quad (21)$$

The objective (15) is the lexicographic minimization of the total amount of displacement by which preferred displacements are missed for improvement requests, followed by the total displacement. The second objective is required because minimizing the sum of the w_m variables does not reduce displacement for non-improvement requests. Constraints (16) and (17) respectively encode the definitions of the decision variables (π_m) and (w_m) described above.

Constraints (18) are added in order to ensure that no airline's total displacement increases significantly from this mechanism. The parameter $\nu \geq 0$ here controls the strictness of this requirement. A higher value of ν gives greater flexibility to satisfy improvement requests, but also means airlines may have to accept a greater schedule displacement. These constraints essentially ensure that the displacement budget mechanism is carried out in a fair manner. It may therefore be natural to ask why we do not use the same fairness constraints (12) introduced in Section 4.3, or similar fairness constraints based on how improvements (w_m) are distributed among the airlines rather than total displacement. The previous fairness metrics were defined in terms of proportion of schedule displacement, as we do not know a priori the exact fair amounts of total schedule displacement that should be allocated to each airline. In the displacement budget mechanism, the reference schedule is by construction a fair baseline, and so considering deviations from this reference is reasonable and leads to a more computationally tractable optimization problem. In addition, the fairness constraints (12) do not prevent increases in displacement for all airlines.

The constraints (19) and (20) ensure that all requests are assigned a displacement cost at most $\max\{\delta_m, \sigma_m\}$; that is, by participating in this mechanism, an airline will receive no more displacement for a request than its baseline displacement, or its preferred displacement if this is greater.

A key feature of the displacement budget mechanism is that it does not require airlines to reveal any potentially commercially sensitive operating costs.

5.2. Constraints on Airline Preferences

In addition to the displacement budget constraint, we impose other constraints on the preferred displacements specified by the airlines. These constraints forbid airlines to specify preferred displacements that are impossible to satisfy exactly. They also prevent airlines assigning excessive amounts of displacements to individual requests. In this way, they not only help to ensure airlines use their displacement budgets more effectively, but also mitigate against them gaming the mechanism.

5.2.1. Displacement Size. Recall that the displacement of a request series m is given by $|\mathcal{D}_m||t - t_m|$, where t is the time interval of the slots allocated to m . In order to satisfy a preferred displacement δ_m exactly, we must therefore have

$$\delta_m \bmod |\mathcal{D}_m| = 0. \quad (22)$$

5.2.2. Turnaround Constraints. The following proposition states that if there is an upper bound on the turnaround time for a request pair, then the difference between preferred displacements for the request series composing this pair is also bounded.

Proposition 3. Suppose $(m_1, m_2) \in \mathcal{P}$ and for the preferred displacements δ_{m_1} and δ_{m_2} , there exist $T_{m_1}, T_{m_2} \in [1, T]$ such that $f_{m_1}^{T_{m_1}} = \delta_{m_1}, f_{m_2}^{T_{m_2}} = \delta_{m_2}$, and $T_{m_2} - T_{m_1} \leq \bar{l}_{m_1, m_2}$. Then it follows that

$$|\delta_{m_1} - \delta_{m_2}| \leq 2|\mathcal{D}_{m_1}|\bar{l}_{m_1, m_2}. \quad (23)$$

Proof. Suppose that there exist T_{m_1} and T_{m_2} satisfying the conditions of the proposition for given $(m_1, m_2) \in \mathcal{P}$. It then follows that

$$\begin{aligned} |T_{m_1} - t_{m_1}| - |T_{m_2} - t_{m_2}| &\leq |T_{m_1} - T_{m_2}| + |t_{m_1} - t_{m_2}| \\ &\leq 2\bar{l}_{m_1, m_2}. \end{aligned}$$

Multiplying through by $|\mathcal{D}_m|$ then yields that

$$|f_{m_1}^{T_{m_1}} - f_{m_2}^{T_{m_2}}| = |\delta_{m_1} - \delta_{m_2}| \leq 2|\mathcal{D}_m|\bar{l}_{m_1, m_2},$$

as required. $\square$

5.2.3. Displacement Upper Bound. It is possible that in order to achieve a reduced delay for some of its requests,

an airline may assign a large portion of its total baseline delay to a few requests that it does not value. If these particular requests are made for off-peak periods, then specifying large preferred displacements may not even lead to the requests being displaced more because the displacement budget mechanism will not increase displacement unnecessarily. To avoid airlines gaming the system in this way, we must impose restrictions on how this displacement is allocated.

If a request pair is made for unsaturated periods (with respect to the reference schedule) and the requested times are consistent with bounds on turnaround time, then the request pair has been allocated its preferred slots. Moreover, because there is already slack capacity at these time periods, displacing this request pair will not immediately help the coordinator satisfy other improvement requests. Therefore, the airline should not be allowed to use its displacement budget on this request pair. More generally, the preferred displacement allocated to a request pair should not be greater than the displacement from the closest feasible pair of slots in unsaturated time periods with respect to the reference schedule.

The minimum displacement from inserting the request pair $(m_1, m_2) \in \mathcal{P}$ into the reference schedule $\bar{x}$ is the optimal value to the following optimization problem. This problem is just the base model restricted to allocating slots to (m_1, m_2) , and where the capacities have been updated using the aggregate reference schedule $\bar{y}$:

$$\text{minimize}_{x_{m_1}^t, x_{m_2}^t} \sum_{t \in \mathcal{T}} f_{m_1}^t x_{m_1}^t + f_{m_2}^t x_{m_2}^t \quad (24)$$

$$\text{subject to } \sum_{t \in \mathcal{T}} x_m^t = 1, \quad m \in \{m_1, m_2\}, \quad (25)$$

$$\begin{aligned} \sum_{t=s}^{s+c-1} (x_{m_1}^t + \bar{y}_{dt}^A) &\leq \alpha_c^{dt}, \quad c \in \mathcal{C}, \\ d &\in \mathcal{D}_{m_1}, s \in [0, T-c], \end{aligned} \quad (26)$$

$$\begin{aligned} \sum_{t=s}^{s+c-1} (x_{m_2}^t + \bar{y}_{dt}^D) &\leq \beta_c^{dt}, \quad c \in \mathcal{C}, \\ d &\in \mathcal{D}_{m_2}, s \in [0, T-c], \end{aligned} \quad (27)$$

$$\begin{aligned} \sum_{t=s}^{s+c-1} (x_{m_1}^t + x_{m_2}^t + \bar{y}_{dt}^A + \bar{y}_{dt}^D) &\leq \gamma_c^{dt}, \\ c &\in \mathcal{C}, d \in \mathcal{D}_1, s \in [0, T-c], \end{aligned} \quad (28)$$

$$L_{m_1, m_2} \leq \sum_{t \in \mathcal{T}} t x_{m_1}^t - \sum_{t \in \mathcal{T}} t x_{m_2}^t \leq \bar{l}_{(m_1, m_2)}, \quad (29)$$

$$x_{m_1}^t, x_{m_2}^t \in \{0, 1\}, \quad \text{for all } t \in \mathcal{T}. \quad (30)$$

We call the above formulation the *minimum insertion problem* and the optimal value the *minimal insertion displacement*. Denoting the minimum insertion displacement to this problem by B_{m_1, m_2} , we then

impose the following constraints on the preferred displacements:

$$\delta_{m_1} + \delta_{m_2} \leq B_{m_1 m_2} \quad \text{for all } (m_1, m_2) \in \mathcal{P}. \quad (31)$$

To find the minimum insertion displacement $B_{m_1 m_2}$, we do not need to solve the minimum insertion problem by integer programming techniques, as it can be calculated efficiently via smart enumeration. Details are given in Online Appendix B.

Whereas (22) and (23) hold immediately for the baseline displacements (i.e., if we set $\delta_m = \sigma_m$), it is not necessarily true that $\sigma_{m_1} + \sigma_{m_2} \leq B_{m_1 m_2}$. The following result gives mild conditions under which this does hold.

Proposition 4. *Suppose that the reference schedule is airline Pareto optimal and that for each $(m_1, m_2) \in \mathcal{P}$ there exists $a \in \mathcal{A}$ such that $m_1, m_2 \in \mathcal{M}_a$. It follows that*

$$\sigma_{m_1} + \sigma_{m_2} \leq B_{m_1 m_2} \quad \text{for each } (m_1, m_2) \in \mathcal{P}.$$

Proof. We suppose that there exist $(m_1, m_2) \in \mathcal{P}$, both requests arising from airline a , for which $\sigma_{m_1} + \sigma_{m_2} > B_{m_1 m_2}$ and obtain a contradiction. Let $T_{m_1}^*, T_{m_2}^*$ be optimal allocated times for the above minimum insertion problem for (m_1, m_2) . Now modify the reference schedule so that requests m_1 and m_2 are allocated slots $T_{m_1}^*$ and $T_{m_2}^*$, respectively. This schedule remains feasible (for the base model) and achieves a lower displacement for airline a . This contradicts the airline Pareto optimality of the reference schedule and concludes the proof. $\square$

Remark 1. The assumption that the reference schedule is airline Pareto optimal is already guaranteed by Algorithm 2. Although it is possible that arrival and departure request series in a request pair can be made by different airlines, this rarely occurs in practice.

5.2.4. Forbidden Time Periods. It is possible that for some time periods, a rolling capacity constraint will have a zero right-hand side. This may occur because of a curfew on flights during the night or when the slot pool is updated during a hierarchical application of the slot-allocation mechanism (see Section 6).

A zero on the right-hand side of an arrival or total movement capacity constraint prevents the allocation of arrival slots for all time periods active in that constraint. This similarly applies to the allocation of departure slots if there is a departure or total capacity constraint with a zero right-hand side.

Denote by $\mathbf{0}$ the *empty schedule*. The set of time periods for which no arrival slots can be allocated consists of those that are arrival-saturated with respect to the empty schedule. We call these *arrival-forbidden periods*. We can similarly define *departure-forbidden periods*. Indicator functions that characterize

arrival and departure forbidden periods, which we denote (F_{dt}^A) and (F_{dt}^D) , respectively, can thus be defined as follows:

$$F_{dt}^A := A_d^t(\mathbf{0}), \quad F_{dt}^D := D_d^t(\mathbf{0}).$$

Because an airline cannot be allocated a slot in a forbidden period, we require that airlines do not specify preferred displacements that imply receiving a slot in such a time period; that is, for an arrival request series $m \in \mathcal{M}^A$, we have

$$\delta_m \neq |D_m|h \text{ if } F_{d,(t_m+h)}^A = 1 \text{ and } F_{d,(t_m-h)}^A = 1 \text{ for any } d \in \mathcal{D}_m, \quad (32)$$

and for a departure request series $m \in \mathcal{M}^D$, we have

$$\delta_m \neq |D_m|h \text{ if } F_{d,(t_m+h)}^D = 1 \text{ and } F_{d,(t_m-h)}^D = 1 \text{ for any } d \in \mathcal{D}_m. \quad (33)$$

5.3. Calculation of Airline Preferences

Airlines are able to specify preferred displacements in an arbitrary way as long as the displacement budget constraint and the constraints in Section 5.2 are satisfied. However, it may not be clear to an airline how best to distribute its displacement in a way that matches its priorities. In this section, we propose a model for distributing an airline's displacement. Although our primary motivation here is to be able to simulate an airline's preferences in order to test the displacement budget mechanism, this model could be used as a decision support tool for the airlines.

The model is formulated and solved for each airline a and assumes that for each $m \in \mathcal{M}_a$ and $t \in \mathcal{T}$, a cost c_m^t is incurred to the airline when request m is allocated slot time t . We call these *airline displacement costs*. Preferences are then calculated by solving the following ILP, which has binary decision variables $(z_m^t)_{m \in \mathcal{M}_a}^{t \in \mathcal{T}}$, which, like the (x_m^t) variables in the slot-allocation model, select a slot time for each request series:

$$\text{minimize } \sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} c_m^t z_m^t \quad (34)$$

$$\text{subject to } \sum_{m \in \mathcal{M}_a} \sum_{t \in \mathcal{T}} f_m^t z_m^t \geq \Sigma_{a'}, \quad (35)$$

$$L_{m_1, m_2} \leq \sum_{t \in \mathcal{T}} t z_{m_1}^t - \sum_{t \in \mathcal{T}} t z_{m_2}^t \leq \bar{l}_{(m_1, m_2)}, \quad (m_1, m_2) \in \mathcal{P}, \quad (36)$$

$$\sum_{t \in \mathcal{T}} f_{m_1}^t z_{m_1}^t + \sum_{t \in \mathcal{T}} f_{m_2}^t z_{m_2}^t \leq B_{m_1 m_2} \quad \text{for all } (m_1, m_2) \in \mathcal{P}, \quad (37)$$

$$z_m^t = 0 \quad \text{for all } m \in \mathcal{M}_a^A \text{ and } t \in \mathcal{T} \text{ such that } F_{dt}^A = 0 \text{ for some } d \in \mathcal{D}_m, \quad (38)$$

$$z_m^t = 0 \quad \text{for all } m \in \mathcal{M}_a^D \text{ and } t \in \mathcal{T} \\ \text{such that } F_{dt}^D = 0 \\ \text{for some } d \in \mathcal{D}_m, \quad (39)$$

$$\sum_{t \in \mathcal{T}} z_m^t = 1 \quad \text{for all } m \in \mathcal{M}_a. \quad (40)$$

$$z_m^t \in \{0, 1\} \quad \text{for all } m \in \mathcal{M}_a, t \in \mathcal{T}.$$

The objective (34) seeks to minimize the total airline displacement costs if all of the airline's displacement preferences are exactly met. The preferred displacements, which will be used as input to the displacement budget mechanism, are then the displacements corresponding to the optimal solution $\mathbf{z}^*$ to the model above, that is,

$$\delta_m = \sum_{t \in \mathcal{T}} f_m^t z_m^{t*}. \quad (41)$$

We now explain how this approach ensures that all constraints on preferred displacements detailed in Section 5.2 hold. The constraint (22) holds directly from the definition of (41), which (along with (40) and the binary constraints) ensures that δ_m is of the form $f_m^t = |\mathcal{D}_m| |t - t_m|$ for some $t \in T$. The turnaround constraints (23) hold by constraints (36) and the application of Proposition 3. The constraint (35) and definition (41) enforce the displacement budget constraint (14). Similarly, (37) and (41) ensure that the displacement upper bound constraint (31) holds. Finally, the constraints (38) and (39) prevent the airline from specifying forbidden periods for their preferred slots and so ensure that (32) and (33) hold.

Given that the above model used airline costs (or disutilities), a natural question concerns why we do not incorporate these into the base model (1)–(9) to directly construct a schedule that incorporates airline preferences. This suggestion is problematic for several reasons. First, these data are likely to be commercially sensitive, and the airlines may not want to reveal them. Second, it would require a consistent definition of cost across all airlines, which may not be possible because of the different operating models of airlines. Third, given that the IATA WSG allows airlines to change details of their requests, such as aircraft type, after the allocation of slots, it may be possible for an airline to game the system by overstating its costs of displacement.

The model above implicitly assumes that the value of a slot for one flight is independent of that for another. In reality, the values of some slots may be linked. For example, for pairs of flights with a potentially large number of connecting passengers, pairs of slots that have compatible connection times will be more valuable than those that do not. Although the development of a model reflecting this is beyond the scope of the current work, we note that existing models in the aviation

literature could be adapted for this purpose. For instance, Pita, Barnhart, and Antunes (2013) proposed an optimization model for airline scheduling that takes into account expected passenger demand and airline cooperation and competition to select slots in a way that maximizes expected profits.

5.3.1. Facilitation of Specification of Airline Preferences.

The displacement budget mechanism requires airlines to specify preferred displacements for potentially many requests subject to several sets of constraints. This process of selecting these preferences could be facilitated by decision support software.

Such software could consist of a graphical interface that allows the user to manually adjust the preferred displacement allocated to each request series (initially the baseline displacements) and an optimization tool that constructs a set of preferred displacements given cost functions for each movement. The graphical interface could also be used to adjust displacements using output from the optimization tool.

The graphical interface should display all request pairs whose displacement can be adjusted from the baseline. For each of these request pairs, the software should display which displacements are feasible. After an adjustment is made, the allowable displacements for other request pairs, which are subject to the displacement budget constraint, should be updated.

The optimization tool could use the model presented in Section 5.3. This requires the user to specify displacement cost functions for each request series. This information would be used only in the optimization and would not be shared with any other parties. To simplify the process of constructing these functions, the optimization tool could make available some parametric forms. For example, in the numerical tests in Section 7, we use airline displacement costs given by

$$c_m^t = R_m |\mathcal{D}_m| |t - t_m|^\eta,$$

where R_m is a weight representing the relative importance of m , and $\eta > 0$ represents the airline's aversion to large displacements.

6. Hierarchical Application of the Slot-Allocation Mechanism

The mechanism described in this paper can take account of priority classes by allocating slots to each group in a separate stage; that is, we allocate slots for the highest priority class first, update the slot pool, then allocate slots to the next highest class, and so on, until slots have been allocated for all priority classes. We call this *hierarchical slot allocation*, and for each priority group, we run the whole slot-allocation mechanism to construct the schedule.

Updating the slot pool consists of updating the capacity constraints. If $\tilde{y}$ is the aggregate schedule for an allocation of slots for the previous priority class, then the right-hand sides of the capacity constraints are updated as follows:

$$\begin{aligned}\alpha_c^{dt} &\leftarrow \alpha_c^{dt} - \sum_{\tilde{s}=s}^{s+c-1} \tilde{y}_{d\tilde{s}}^A, \\ \beta_c^{dt} &\leftarrow \beta_c^{dt} - \sum_{\tilde{s}=s}^{s+c-1} \tilde{y}_{d\tilde{s}}^D, \\ \gamma_c^{dt} &\leftarrow \gamma_c^{dt} - \sum_{\tilde{s}=s}^{s+c-1} (\tilde{y}_{d\tilde{s}}^A + \tilde{y}_{d\tilde{s}}^D).\end{aligned}$$

The lexicographic approach to incorporate priority classes into slot allocation, as deployed by Ribeiro et al. (2018) and described in Section 1.2, cannot be used within our mechanism because we use a two-stage approach to allocate slots to each priority level rather than solve a single optimization problem.

7. Numerical Tests

In this section, we test our slot-allocation mechanism using data for a coordinated airport. The purpose of these tests is to demonstrate how the mechanism performs for a realistic slot-allocation problem, and to also investigate how this performance depends on the different parameters of the mechanism. Testing performance while varying mechanism parameters is problematic when solving the problem hierarchically. This is because when the problem is solved for lower-priority groups, as well as depending on the mechanism parameters, performance will also depend on the slot pool, and these effects cannot easily be disentangled. In order to experiment with the mechanism settings, we therefore solve the problem nonhierarchically. Recall that solving the problem nonhierarchically means disregarding the priority groups and allocating slots to all requests simultaneously. Using the best settings found in these experiments, we will then demonstrate the mechanism using the full hierarchical approach.

This section is organized as follows. In Section 7.1 we describe our general experimental set-up. In Section 7.2 we present our experiments for the non-hierarchical approach, and in Section 7.3 we demonstrate our mechanism using the hierarchical approach. Finally in Section 7.4 we discuss and compare the results of these two approaches.

7.1. Setup

7.1.1. Problem Data. We solve the slot-allocation problem for a coordinated airport. The rolling capacity parameters for this problem are given in Table 2. Because the shortest duration of a rolling capacity constraint is 15 minutes, we use this as the length of

Table 2. Rolling Capacity Constraints for a Coordinated Airport

Arrivals	Departures	Total	
60 minutes	60 minutes	15 minutes	60 minutes
8	12	5	20

our coordination time intervals. There are thus 96 time intervals in each day, and because there are 212 days in this scheduling season, there are a total of 20,352 time periods in the scheduling season.

The number of requests series and individual requests for each priority group is shown in Table 3. Note that because of a lack of historical request times for the “changes to historic” requests, we amalgamate the “historic” and “changes to historic” requests into a single priority group.

The required turnaround time for an aircraft depends on the aircraft model, and airport services, but in general a turnaround time of one hour is sufficient for most aircraft. However, there are request pairs in our data set where the requested arrival and departure times are separated by less than this period. In order to use demand-based fairness, Proposition 2 requires that the requested slot times of request pairs are compatible with minimum turnaround times. We therefore use the following formula to set the minimum turnaround time:

$$l_{m_1 m_2} = \min\{4, t_{m_2} - t_{m_1}\},$$

where four coordination time intervals correspond to one hour. Because it is not strictly necessary, and we do not have the required data, we do not enforce upper bounds on turnaround time.

7.1.2. Airline Displacement Costs. Airline displacement cost data are not available, and therefore they have been constructed synthetically. For the purposes of these numerical tests, we assume that the airline displacement costs are given by

$$c_m^t = R_m |\mathcal{D}_m| |t - t_m|^\eta.$$

This is an increasing function with respect to deviation from the requested slot time. A superlinear function is used to represent the airlines’ costs to reflect the fact that airlines are much more sensitive to

Table 3. Request Data

Priority class	Request series	Requests
Historics	748	14,034
New entrants	76	2,680
Others	1,290	16,328
Total	2,114	33,042

larger displacements. A superlinear cost function was also used in the slot-allocation model in the paper by Castelli, Pellegrini, and Pesenti (2012).

To capture the relative importance of each request, this function is weighted by $R_m|\mathcal{D}_m|$, where R_m is the normalized passenger capacity of the aircraft. In particular, $R_m = \frac{N_m}{456}$, where N_m is the number of seats of the aircraft for which the request is made, and 456 is the maximum aircraft capacity over all requests. Passenger-based costs were used for a similar purpose in the context of a ground-holding problem in Vossen and Ball (2006).

7.1.3. Computer and Solver Setup. The integer linear programs are solved using Gurobi 8.10 (Gurobi Optimization 2018) on a desktop computer that has an Intel(R) Core(TM) i5-4690 processor with four cores. Many of the rolling capacity constraints (5)–(7) will be redundant and binding only for periods of high slot demand. Therefore, to solve the problem more efficiently, we use a branch-and-cut scheme, where the rolling capacity constraints are added as lazy constraints through a solver callback. To do this, the Gurobi LazyConstraints parameter is set to true.

The most computationally demanding part of the slot-allocation mechanism is the calculation of the efficient frontier. The use of fairness constraints can increase the required solution time by orders of magnitude. Whereas the base model for the nonhierarchical problem takes around 5 to 10 minutes to solve, the fairness model can take hours, and is sometimes not solved at all within a reasonable time.

We therefore limit the solution time for each problem to two hours using the Gurobi TimeLimit parameter.

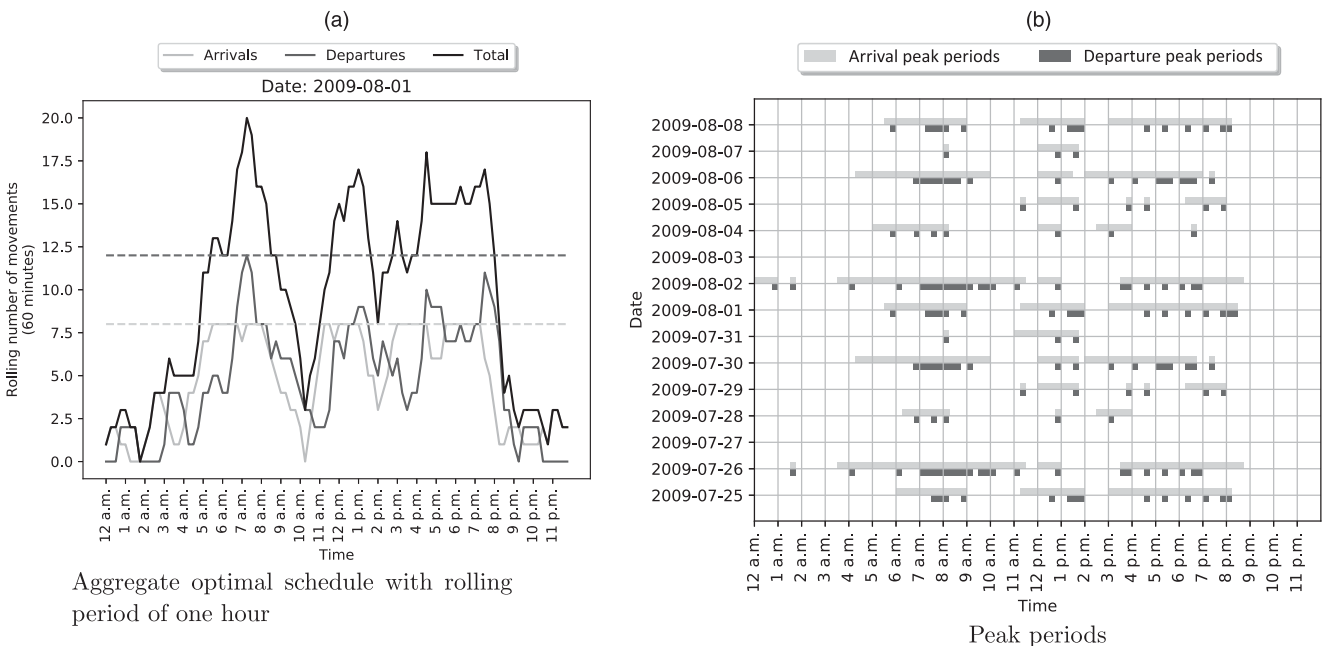
As described in Section 4.5, we use the α -acceptable price of fairness criterion to select the fair reference schedule. Using this approach, we can terminate our ϵ -constraint algorithm once the price of fairness has been exceeded. However, we need not solve an ϵ -fairness problem to optimality to verify that it yields a schedule that exceeds the acceptable price of fairness. The solution algorithm can be terminated once the current lower bound on the optimal value exceeds the acceptable price of fairness. We do this by setting the Gurobi parameter BestBdStop appropriately. All other Gurobi parameters are left with their default values.

7.2. Nonhierarchical Results

In this section, we allocate slots using a nonhierarchical approach; that is, we allocate all slots simultaneously, disregarding priority classes.

7.2.1. Peak Periods and Fairness. To calculate demand-based fairness, we first need to solve the base model and calculate the peak periods. Figure 4(a) shows the aggregate hourly schedule of an optimal solution for the base model on the busiest day in the scheduling season. For a specified time, this gives the numbers of arrivals, departures, and total movements that take place over the next hour. The dashed lines represent the hourly capacities for each type of constraint. The arrival and departure peak periods corresponding to this optimal schedule are shown in Figure 4(b) for a 15-day period including the busiest day. This figure shows that departure peak periods

Figure 4. Rolling Aggregate Schedule and Peak Periods for the Nonhierarchical Problem



are more intermittent than arrival peak periods, which is due to the fact that arrival capacity limits are significantly lower than departure capacity limits. This difference in peak periods for arrivals and departures emphasizes the importance of treating these two types of requests separately.

In Figure 5(a), we plot the total displacement–MDA efficient frontiers. Although we use a reference schedule constructed using demand-based fairness, for the purposes of comparison, we also calculate the efficient frontier for non-demand-based fairness. The points encircled by black lines correspond to the schedules that would be selected under the acceptable price of fairness rule. We use a decrement rate of $\delta = 0.8$ for these calculations. Note that in this figure, we plot the actual values of MDA found by solving these problems, rather than the values of ϵ .

Recall that the smaller the value of the MDA metric, the fairer the schedule. Figure 5(a) shows that the trade-off for displacement versus fairness is better for demand-based fairness; that is, we can construct fairer schedules with smaller values of total displacement. This suggests that for demand-based fairness, it is easier to redistribute schedule displacement to airlines with a high proportion of peak requests rather than to airlines with a high proportion of total requests.

The solution of the ϵ -fairness problems is by far the most computationally expensive component of the proposed slot-allocation mechanism. For both demand-based and non-demand-based fairness, the ϵ -constraint algorithm terminated because the time limit, rather than the acceptable price of fairness, was exceeded. In Figure 5(b), we plot the solution time for each ϵ -fairness problem. This shows that the tighter the fairness constraints, the more difficult

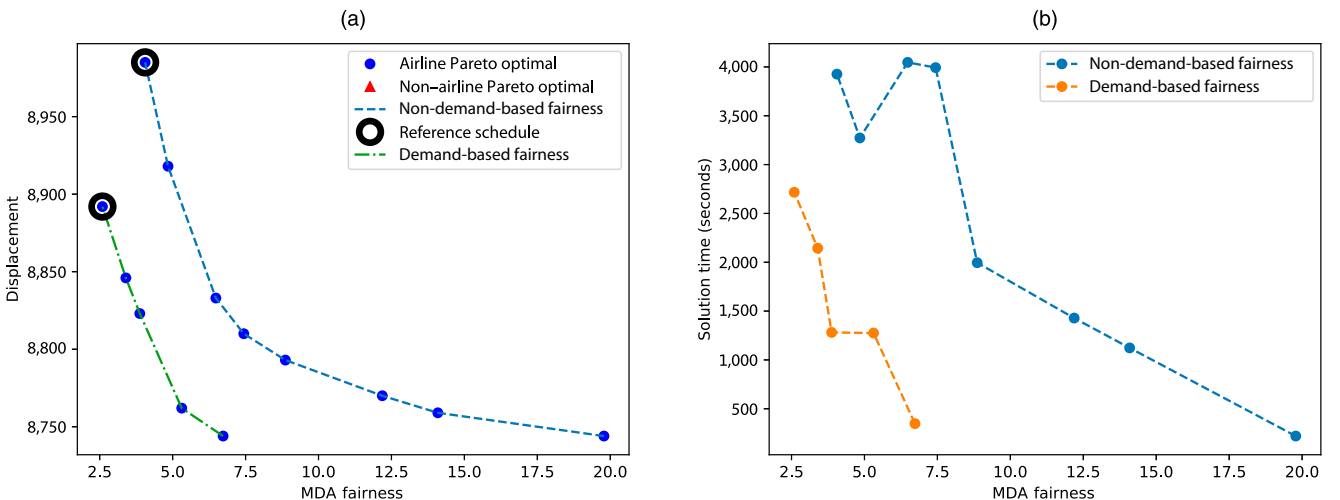
the problem becomes to solve. This suggests that the acceptable price of fairness criterion is therefore quite appropriate, as it massively reduces the required computation time to calculate the fair reference schedule. This plot also shows that the non-demand-based fairness problem is more difficult to solve than the demand-based version. Again, this may be because it is easier to redistribute displacement to airlines with a high proportion of peak requests.

7.2.2. Displacement Budget Mechanism. We now study the performance of the displacement budget mechanism. We denote by $\bar{x}$ and $\bar{y}$, respectively, the reference schedule and aggregate reference schedule, and by $\tilde{x}$ and $\tilde{w}$ the solutions returned from the displacement budget formulation in (15)–(21). Because the displacement budget mechanism takes only a few seconds to run, we do not present the run times.

Distribution of Displacement Adjustments. We first study the distribution of schedule adjustments. Although we present results only when the mechanism is run for parameters $\nu = 0.1$ and $\eta = 1.8$, we have observed the same trends when the mechanism is run under different settings. The results are plotted in Figure 6. This shows the number of requests for every possible combination of requested adjustments ($\frac{\sigma_m - \delta_m}{|\mathcal{Q}_m|}$) and actual adjustments ($\frac{\sum_{i \in \mathcal{A}} f_m^i x_m^i - \sigma_m}{|\mathcal{Q}_m|}$).

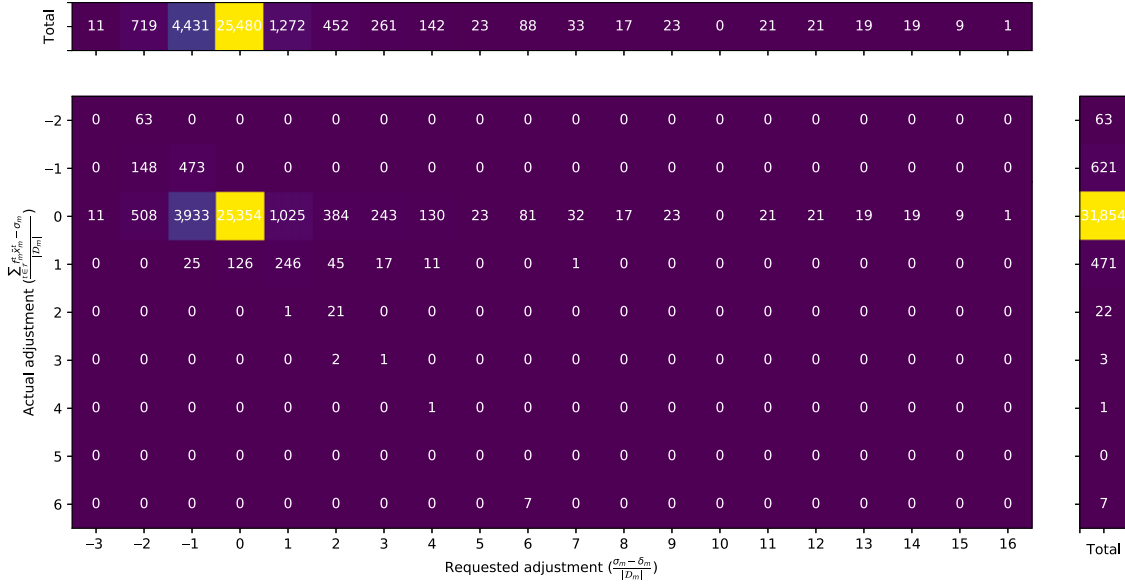
This shows that for the majority of requests (around 75%), the airlines ask for and receive no adjustment to their baseline displacement. We also observe that the range of actual adjustments is much narrower than the range of demanded adjustments. Although a significant number of requests are made for large improvements, up to 16 coordination time intervals

Figure 5. (Color online) Comparison of Demand-Based and Non-Demand-Based Fairness Schemes



Displacement-MDA fairness efficient frontiers

Solution times for ϵ -fairness problems

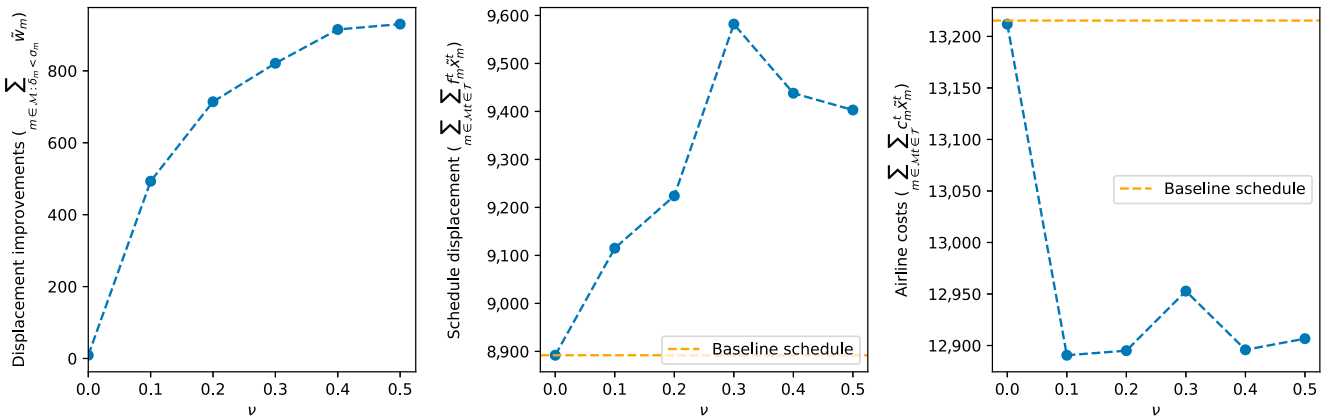
Figure 6. (Color online) Distribution of Adjustments from Displacement Budget Mechanism

(4 hours), no request has its displacement reduced by more 6 coordination time intervals (1.5 hours), and the majority of actual improvements are less than or equal to 2 coordination time intervals (0.5 hours). This suggests that a better strategy for airlines participating in the displacement budget mechanism is to demand smaller adjustments.

We observe that although there are many requests whose displacement can be increased (those with negative requested adjustments), only a fraction (around 11%) of these actually receive extra displacement. Interestingly, some of the requests receive a greater reduction in displacement than demanded, including even a small number of requests for which no or a negative adjustment was asked for. This is because when the mechanism increases displacement for some requests in order to satisfy improvement requests, more capacity is created than necessary, and the displacement of other requests can be reduced.

Sensitivity Analysis with Respect to ν . We now study the performance of the methodology with respect to the parameter ν , which controls the proportion of extra displacement each airline can be allocated. We run the displacement budget mechanism for $\nu = 0.0, 0.1, \dots, 0.5$ with $\eta = 1.8$ as in the previous test. We look at three different attributes of performance. The first is the total improvements, $\sum_{m \in \mathcal{M}: \delta_m < \sigma_m} \tilde{w}_m$; the second is the schedule displacement, $\sum_{m \in \mathcal{M}} \sum_{t \in \mathcal{T}} f_m^t \tilde{x}_m^t$; and the third is the total airline cost, $\sum_{m \in \mathcal{M}} \sum_{t \in \mathcal{T}} c_m^t \tilde{x}_m^t$. The results are shown in Figure 7.

This shows that when no increase in displacement is allowed ($\nu = 0$), no schedule improvements can be made. Because no improvements are made for $\nu = 0$, the schedule displacement and airline costs do not change from the baseline. As ν increases, the total amount of schedule improvements increases. Interestingly, although the schedule displacement increases initially with ν , for larger ν , it then decreases.

Figure 7. (Color online) Performance of Displacement Budget Mechanism for Different Values of ν 

This means that for higher values of ν , increases in displacement are being redistributed to fewer airlines. This unfair redistribution of displacement for higher ν could mean that the size of the largest displacements for some airlines is increasing. Because large displacements contribute more to airline costs, this may explain why airline costs seem to flatline for higher values of ν even though the total schedule improvements are increasing.

Sensitivity Analysis with Respect to η . The parameter η controls how averse airlines are to larger displacements: the larger the value of η , the more averse they are. Although in reality a coordinator would have no control over η , some useful insights can be gained about how the mechanism performs with respect to the airlines' preferences by varying this parameter. In this experiment, we run the displacement budget mechanism for $\eta = 0.0, 0.2, \dots, 2.6$, with fixed $\nu = 0.1$. For each value of η , we plot the relative reduction in total airline costs with respect to the reference schedule, that is,

$$\frac{\sum_{m \in \mathcal{M}} \sum_{t \in \mathcal{T}} c_m^t \tilde{x}_m^t - \sum_{m \in \mathcal{M}} \sum_{t \in \mathcal{T}} c_m^t \tilde{x}_m^t}{\sum_{m \in \mathcal{M}} \sum_{t \in \mathcal{T}} c_m^t \tilde{x}_m^t},$$

where, recall, $c_m^t = |\mathcal{D}_m| R_m |t - t_m|^\eta$. The results are plotted in Figure 8.

This figure shows that for small values of η , the displacement budget mechanism results in a schedule with greater total airline costs. This is because the displacement budget mechanism increases the overall schedule displacement, and for small values of η , the total airline costs and schedule displacement as functions have a similar shape. For higher values of η , the reduction in airline costs increases, and reaches around 3% for $\eta = 2.6$. This is likely due to the mechanism reducing the largest displacements, which for

high values of η , contribute much more to the total airline costs than small displacements. These results suggest that the more averse the airlines are to large displacements, the greater the benefit of the displacement budget mechanism.

7.3. Hierarchical Results

For this experiment we solve the slot-allocation problem using the hierarchical approach described in Section 6. The peak periods and efficient frontiers for each priority group are shown in Figure 9. In the cases of historical and new entrant requests, the ϵ -constraint algorithm terminates when the acceptable price of fairness is exceeded, whereas for the other requests it terminates because of the time limit. A summary of the results of the displacement budget mechanism at each stage is shown in Table 4.

7.3.1. Historic Requests. From Figure 9(a), we can see that the peak periods seem to be largely confined to Tuesday and Friday mornings. By solving the efficient frontier, we were able to improve the MDA from around 6.0 (by solving the base model) to 1.0, after which the price of fairness is exceeded.

In the displacement budget mechanism, although there is a total of 243 individual requests for improvement, no improvements are made. This is to be expected, as only a very small number of the historical requests receive any displacement, which means that the opportunities for slot exchange are very small. There is a very small change in the total airline costs due to the displacement budget mechanism returning a different schedule from the reference schedule.

7.3.2. New Entrants. There are relatively few new entrants request series (76 out of 2,114), and so the peak periods for these requests largely coincide with the peak periods for the historic request series. Because the slots for the historical request series have already been allocated at this point, these periods do not have any slack capacity and are therefore considered to be forbidden for the purposes of the displacement budget mechanism. Solving the base model yields a schedule that is already very fair, with an MDA of around 1.0. Although fairer schedules can be found, none of these are airline Pareto optimal. For the displacement budget mechanism, there are 28 individual improvement requests, and the mechanism is able to make 27 total displacement improvements. This is done without increasing the schedule displacement for new entrants, and the new schedule has significantly lower airline costs.

7.3.3. Others. The peak periods for the other requests resemble those that were calculated for the solution of the nonhierarchical problem, but here many of the

Figure 8. (Color online) Relative Reduction in Total Airline Costs for Different Values of η

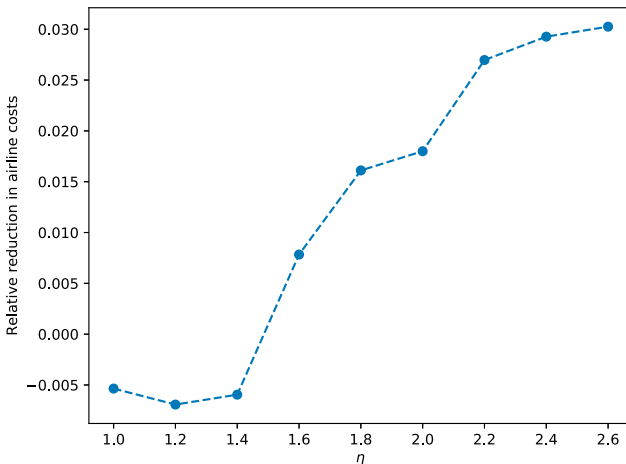
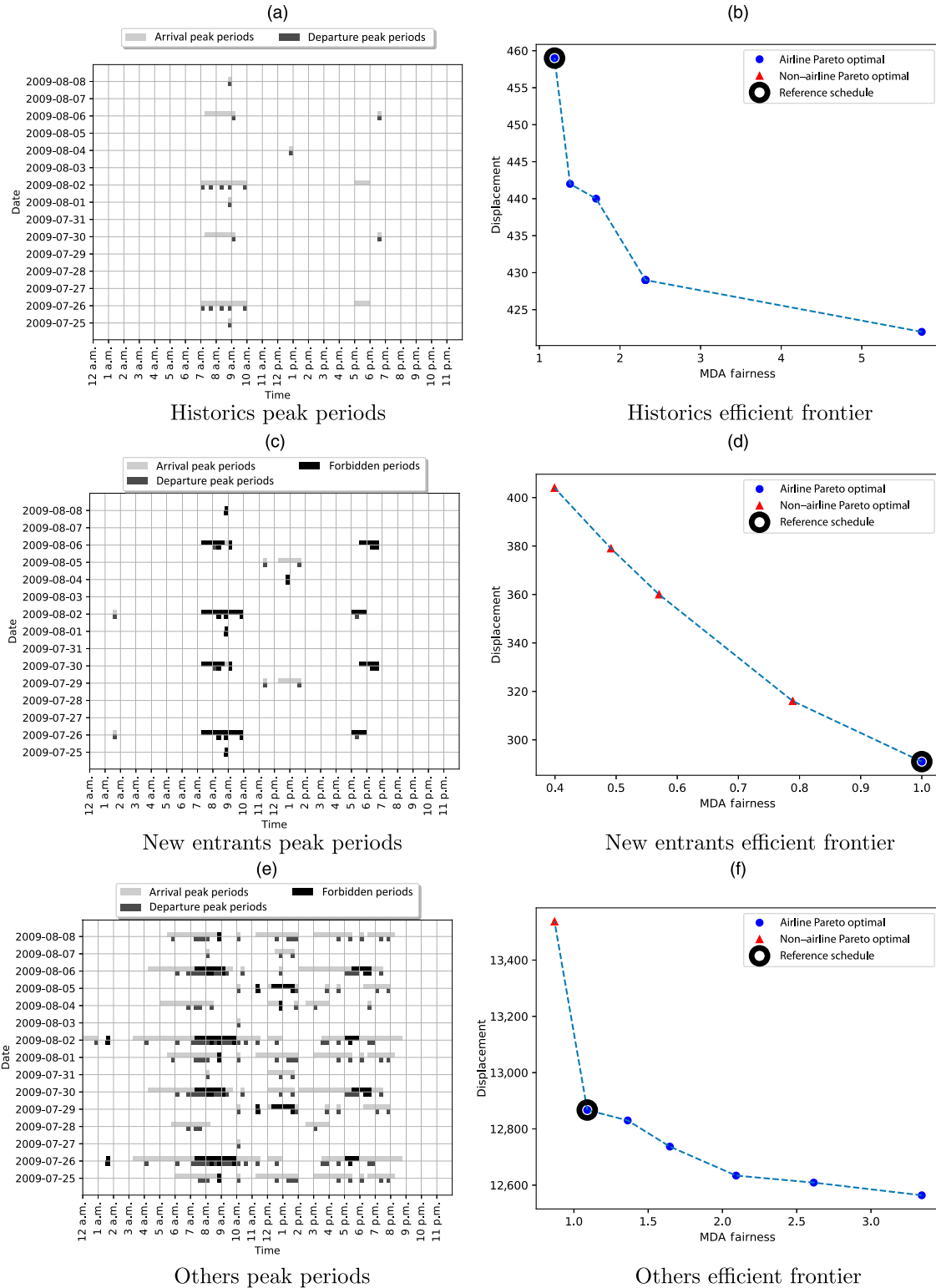


Figure 9. (Color online) Peaks Periods and Efficient Frontiers for Hierarchical Slot Allocation

periods are forbidden. For the calculation of the efficient frontier, the value of the MDA can be improved from an initial MDA value of around 3.3 to 1.1, while the solutions remain airline Pareto optimal. In the

displacement budget mechanism, a total displacement improvement of 384 over 2,128 individual improvement requests, which leads to an improvement in airline costs of 1.2%.

Table 4. Summary of Displacement Budget Mechanism for Hierarchical Slot Allocation

	Historics	New entrants	Others
Baseline schedule displacement	459	291	12,867
New schedule displacement	459	291	13,174
Number of improvement requests	243	28	2,128
Total displacement improvements	0	27	384
Baseline airline costs	335.80	138.39	18,395.04
New airline costs	336.39	98.45	18,179.16

7.4. Discussion

We note that for both the nonhierarchical and hierarchical approaches, the fairness of a schedule can in almost all cases be greatly improved with only small increases in the total displacement of a schedule. Most ϵ -fairness problems that were solved yielded an airline-Pareto-optimal schedule. The only exceptions to this were schedules with an MDA of less than 1. Inspection of airline-Pareto-optimal solutions with an MDA near 1 reveals that there are a number of airlines with a very small proportion of peak requests, but that have zero displacement, which means the deviation from absolute fairness of these airlines' fairness indices is 1. When $\epsilon < 1$, these airlines must be assigned some displacement, and the solutions become non-airline Pareto optimal. We surmise that this is because, although these airlines have some requests during peak periods, no request in these periods needs to be displaced. This further motivates the need for the development of a more sophisticated measure of peak periods, as discussed at the end of Section 4.2.

In terms of schedule improvements and reduced airline costs, the displacement budget seems to be slightly more effective for the nonhierarchical than the hierarchical approach. This is to be expected, as allocating slots to higher-priority groups greatly restricts the slots that can be allocated to lower-priority groups, and also allocating slots to smaller groups of requests reduces the opportunities for exchanging slots. This reasoning suggests that the mechanism may be more effective for larger and more congested airports, where there would be more opportunities for slot exchange. Also, if we could relax the requirement for all requests in a request series to be allocated slots at the same time, we would expect the overall performance of the displacement budget mechanism to improve as a result of the greater flexibility in scheduling.

8. Conclusions

In this paper, we propose a mechanism for allocating slots at congested airports that incorporates efficiency, fairness, and airline preferences. The mechanism consists of constructing a fair reference schedule and then adjusting this to better match the

participating airlines' priorities. The construction of the fair reference schedule uses a new fairness measure that was developed on the principle that the proportion of schedule displacement an airline receives should depend on the number of requests it makes during peak demand periods. The approach was tested using request data from a coordinated airport. Our fairness approach outperformed the previous fairness metric of Zografos and Jiang (2019) in terms of trade-off with schedule displacement and required computation time. The displacement budget mechanism was shown to be able to make a significant number of improvements requested by airlines without large increases in schedule displacement.

There are numerous possible extensions to our approach. First, the base model could be modified to model the IATA WSG to a higher fidelity. For example, although the airport on which we tested our model had more than sufficient capacity to allocate slots to all requests, in general, this will not be true. Therefore, the model should be extended to allow some requests to be rejected (as is done in Ribeiro et al. 2018). In the case where there are likely to be a lot of rejected requests, there may also be a need to enforce the fair distribution of rejections to airlines in the model.

Another important extension to this work would be the development of a peak indicator that not only indicates whether a period has peak demand, but also measures the severity of that demand. This would allow one to distinguish between requests made in periods with different levels of demand. As discussed at Section 7.4, such an indicator may lead to improved fairness outcomes. Such an indicator could also be used to improve the displacement budget mechanism, by making it more expensive for airlines to reduce displacement during highly peak periods. By encouraging airlines to make more attainable requests, more improvement requests could be satisfied.

Finally, there is a need to develop faster algorithms for solving our slot-allocation problems, particularly for the ϵ -fairness problem, whose solution is the most time-consuming part of the proposed slot-allocation mechanism. This will be necessary in order to run the mechanism for larger airports.

Acknowledgments

This work was supported by the Engineering and Physical Sciences Research Council [Grant EP/M020258/1]. The authors thank Rob Shone for his detailed feedback on a draft version of this manuscript. Because of confidentiality concerns, supporting data cannot be made openly available. Additional details relating to the data are available from the Lancaster University data archive at <https://dx.doi.org/10.17635/lancaster/researchdata/295>. The opinions expressed in this paper reflect the authors' views.

Endnotes

¹Note that our use of the term will differ slightly from that in the International Air Transport Association (2017) WSG. There, the term slot pool is used for the slots remaining after all historic slots have been allocated, whereas we use it for any stage of the allocation.

²This is usually defined as the total benefit to all participants.

³This is an initiative used in the United States to hold aircraft on the ground to avoid airborne delays.

References

- Ball M, Donohue GL, Hoffman K (2005) Auctions for the safe, efficient, and equitable allocation of airspace system resources. Cramton P, Shoham Y, Steinberg R, eds. *Combinatorial Auctions* (MIT Press, Cambridge, MA), 507–538.
- Barnhart C, Bertsimas D, Caramanis C, Fearing D (2012) Equitable and efficient coordination in traffic flow management. *Transportation Sci.* 46(2):262–280.
- Benlic U (2018) Heuristic search for allocation of slots at network level. *Transportation Res. Part C: Emerging Tech.* 86(January):488–509.
- Bertsimas D, Gupta S (2015) Fairness and collaboration in network air traffic flow management: an optimization approach. *Transportation Sci.* 50(1):57–76.
- Bertsimas D, Farias VF, Trichakis N (2012) On the efficiency-fairness trade-off. *Management Sci.* 58(12):2234–2250.
- Castelli L, Pellegrini P, Pesenti R (2012) Airport slot allocation in Europe: Economic efficiency and fairness. *Internat. J. Revenue Management* 6(1–2):28–44.
- Corolli L, Lulli G, Ntamo L (2014) The time slot allocation problem under uncertain capacity. *Transportation Res. Part C: Emerging Tech.* 46(September):16–29.
- Council of the European Union (1993a) Council regulation (EEC) no. 95/93 of 18 January 1993 on common rules for the allocation of slots at community airports. Accessed August 1, 2019, <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:31993R0095>.
- Council of the European Union (1993b) Regulation (EC) no. 793/2004 of the European Parliament and of the council of 21 April 2004 amending council regulation (EEC) no. 95/93 on common rules for the allocation of slots at community airports. Accessed August 1, 2019, <https://eur-lex.europa.eu/legal-content/en/ALL/?uri=CELEX:32004R0793>.
- Dal Sasso V, Fomeni FD, Lulli G, Zografos KG (2019) Planning efficient 4D trajectories in air traffic flow management. *Eur. J. Oper. Res.* 276(2):676–687.
- Evans A, Vaze V, Barnhart C (2016) Airline-driven performance-based air traffic management: Game theoretic models and multicriteria evaluation. *Transportation Sci.* 50(1):180–203.
- Fairbrother J, Zografos KG (2018) On the development of a fair and efficient slot scheduling mechanism at congested airports. Presentation, 97th Transportation Research Board Annual Meeting, January 7–11, Transportation Research Board, Washington, DC.
- Gurobi Optimization (2018) Gurobi Optimizer reference manual, version 8.1. Accessed August 1, 2019, <https://www.gurobi.com/>.
- International Air Transport Association (2017) Worldwide slot guidelines. Accessed August 1, 2019, <https://www.iata.org/policy/slots/Documents/wsg-8-english.pdf>.
- Jacquillat A, Odoni A (2015) An integrated scheduling and operations approach to airport congestion mitigation. *Oper. Res.* 63(6):1390–1410.
- Jacquillat A, Vaze V (2018) Interairline equity in airport scheduling interventions. *Transportation Sci.* 52(4):941–964.
- Kalai E, Smorodinsky M (1975) Other solutions to Nash's bargaining problem. *Econometrica* 43(3):513–518.
- Morrison S, Winston C (2007) Another look at airport congestion pricing. *Amer. Econom. Rev.* 97(5):1970–1977.
- Nash J (1950) The bargaining problem. *Econometrica* 18(2):155–162.
- Pellegrini P, Castelli L, Pesenti R (2012) Secondary trading of airport slots as a combinatorial exchange. *Transportation Res. Part E: Logist. Transportation Rev.* 48(5):1009–1022.
- Pellegrini P, Bolić T, Castelli L, Pesenti R (2017) SOSTA: An effective model for the simultaneous optimisation of airport slot allocation. *Transportation Res. Part E: Logist. Transportation Rev.* 99(March):34–53.
- Pilon N, Cook A, Ruiz S, Bujor A, Castelli L (2016) Improved flexibility and equity for airspace users during demand-capacity imbalance. Schaefer D, ed. *6th SESAR Innovation Days, November 8–10, Delft, Netherlands*.
- Pita JP, Barnhart C, Antunes AP (2013) Integrated flight scheduling and fleet assignment under airport congestion. *Transportation Sci.* 47(4):477–492.
- Rassenti S, Smith V, Bulfin R (1982) A combinatorial auction mechanism for airport time slot allocation. *Bell J. Econom.* 13(2):402–417.
- Ribeiro NA, Jacquillat A, Antunes A, Odoni A, Pita J (2018) An optimization approach for airport slot allocation under IATA guidelines. *Transportation Res. Part B: Methodological* 112(June):132–156.
- Swaroop P, Ball MO (2013) Consensus-building mechanism for setting service expectations in air traffic flow management. *Transportation Res. Record* 2325(1):87–96.
- Vaze V, Barnhart C (2012) Airline frequency competition in airport congestion pricing. *Transportation Res. Record* 2266(1):69–77.
- Vossen T, Ball MO (2006) Slot trading opportunities in collaborative ground delay programs. *Transportation Sci.* 40(1):29–43.
- Vossen T, Ball M, Hoffman R, Wambganss M (2003) A general approach to equity in traffic flow management and its application to mitigating exemption bias in ground delay programs. *Air Traffic Control Quart.* 11(4):277–292.
- Wambganss M (1996) Collaborative decision making through dynamic information transfer. *Air Traffic Control Quart.* 4(2):107–123.
- Zografos K, Jiang Y (2016) Modelling and solving the airport slot scheduling problem with efficiency, fairness, and accessibility considerations. *9th Triennial Sympos. Transportation Anal. (TRISTAN)*, June 13–17, Oranjestad, Aruba.
- Zografos K, Jiang Y (2017) Modelling fairness in slot scheduling decisions at capacity-constrained airports. Presentation, 96th Transportation Research Board Annual Meeting, January 8–12, Transportation Research Board, Washington, DC.
- Zografos K, Jiang Y (2019) A bi-objective efficiency-fairness model for scheduling slots at congested airports. *Transportation Res. Part C: Emerging Tech.* 102(May):336–350.
- Zografos KG, Androutsopoulos KN, Madas MA (2018) Minding the gap: Optimizing airport schedule displacement and acceptability. *Transportation Res. Part A: Policy Practice* 114(August):203–221.
- Zografos KG, Salouras Y, Madas MA (2012) Dealing with the efficient allocation of scarce resources at congested airports. *Transportation Res. Part C: Emerging Tech.* 21(1):244–256.