

## A CONTAINER SHIP TRAFFIC MODEL FOR SIMULATION STUDIES

JAKUB WAWRZYNIAK <sup>a</sup>, MACIEJ DROZDOWSKI <sup>a,\*</sup>, ÉRIC SANLAVILLE <sup>b</sup>

<sup>a</sup>Institute of Computing Science  
Poznań University of Technology  
Piotrowo 2, 60-965 Poznań, Poland  
e-mail: Maciej.Drozdzowski@cs.put.poznan.pl

<sup>b</sup>Faculty of Science and Technology  
Normandy University—UNIHAVRE, UNIROUEN, INSA Rouen, LITIS  
25 rue Philippe Lebon, BP 540 76086 Le Havre, France

The aim of this paper is to develop a container ship traffic model for port simulation studies. Such a model is essential for terminal design analyses and testing the performance of optimization algorithms. This kind of studies requires accurate information about the ship stream to build test scenarios and benchmark instances. A statistical model of ship traffic is developed on the basis of container ship arrivals in eight world ports. The model provides three parameters of the arriving ships: ship size, arrival time and service time. The stream of ships is divided into classes according to vessel sizes. For each class, service time distributions and mixes of return time distributions are provided. A model of aperiodic arrivals is also proposed. Moreover, the results achieved are used to compare port specific features.

**Keywords:** benchmark instances, algorithm evaluation, data analysis, scheduling, container ship traffic modeling, marine logistics.

### 1. Introduction

Ports and maritime container terminals are vital infrastructure to the global economy. Due to competition and increasing traffic, they are intensively optimized. There are many operations research formulations optimizing elements of port logistics: berth assignment problem, tugboat and quay crane assignment problems, landside container traffic routing (Bierwirth and Meisel, 2010; 2015; Buhrkal *et al.*, 2011; Chen and Yang, 2014; Çagatay and Siu Lee Lam, 2021; Giallombardo *et al.*, 2010; Hedjar and Bounkhe, 2019; Imai *et al.*, 2014; Kang *et al.*, 2020; Liu, 2020; Stahlbock and Voß, 2008; Wawrzyniak *et al.*, 2020; Wang *et al.*, 2020). All these problems are solved with some competing combinatorial optimization algorithms.

Yet, the test instances are often built without considering the actual data features. Classic operations research and combinatorial optimization problems take advantage of existing benchmark data sets. For example, there are benchmarks for shop scheduling

problems (Taillard, 1993), Parallel Workload Archives for supercomputer jobs scheduling (Feitelson *et al.*, 2014), NEO instances for vehicle routing problem (NEO Research Group, 2013), and even on a broader scope TSPLIB (Reinelt, 1995), OR-Library (Beasley, 2018) benchmark collections. The above-mentioned problems of port logistics also need benchmark instances to test algorithm runtime and their solutions quality. Until recently there was limited availability of such data.

The aim of this paper is to develop a model of container ship traffic. This model can then be used to build input data for the benefit of algorithm performance analyses in the above optimization problems of port logistics. Since the model is built using historical data from eight world ports, the test instances have realistic features and, therefore, are better suited than *ad hoc* constructed data sets.

A ship traffic model (STM) is also essential in simulation studies when a new container terminal is designed or an existing one is redesigned. In many port optimization formulations, the partitioning of the quay into berths, the layout of the storage yard and

\*Corresponding author

the landside container transport lanes are considered as given. These design decisions further impact future ship to berth assignments, tugboats and crane assignments, container storage and transport organization. Since container terminals are complex infrastructures whose operations evolved over generations, there are many material, environmental, legal, and financial constraints determining ship berths, crane schedules, vehicle logistics, container storage, and other port activities. Due to this complexity, it is hard to expect that analytical models will provide accurate estimations of the effects of the planned port changes. Furthermore, large time horizon of the future terminal use implies large uncertainty, which is difficult to take into account in analytical models. Therefore, simulation studies are essential in assessing future port throughput and performance for a given terminal layout. Such studies require an accurate but versatile ship traffic model providing the right inputs for the simulation. In this paper we show how to build such a model.

A realistic ship traffic model provides indications on realities of the assumptions often made when studying design and optimization problems in port logistics. Thus, it allows to avoid *ad hoc* choices in studying such problems. For example, we established that: (i) ship length is the key determinant for return and service times (hence these three parameters are strongly related), (ii) return times have a strong periodic component (and memoryless distributions are *unjustified*), (iii) typical distributions (uniform, normal, exponential) are the least suitable to represent service times, (iv) aperiodic arrivals have a strong seasonal component, (v) ship return time patterns determining time horizons of planning and scheduling easily span over years. Although the mathematical structure of the model is the same for all ports, its instantiations differ. As a side benefit, this gives also a way to compare ports.

The ship traffic model introduced in this paper is explainable. That is, the model not only generates ship streams with given characteristics, but it can also help to explain how and why these characteristics emerge. Thus, rather than some machine-learning “black-box” (cf. Gosasang *et al.* (2011)), we use data analysis and explicit statistical distributions. As a result, the model parameters can be used directly in algorithms solving certain port logistic problems. Let us note that our ship traffic model is not predictive in the classical sense in which, e.g., time series analysis is predictive. We do not intend to predict the future number of TEUs or calling ships on the basis of some independent determinants. The model is built to take advantage of recreating features of the real traffic in the above-mentioned simulation applications. The number of calling ships is an input setting. It is assumed that according to the current state of the art, numbers from determined probability distributions can be

generated pseudo-randomly with satisfactory accuracy.

The organization of the text is the following. In the next section, we shortly outline related literature. The rationale behind the ship traffic model is presented in Section 3. Section 4 is dedicated to the method of ship traffic model construction. In Section 5, we outline how to use the STM to generate ship stream data. In Section 6 ports are compared using relationships in the analyzed data. The last section concludes this work. The notation is summarized in Table 1. Due to the size of the collected data, only selected example results are presented in this text. We invite the readers to refer to the work of Wawrzyniak *et al.* (2021), where details of the built ship traffic model are collected.<sup>1</sup>

## 2. Related work

Most current models of vessel arrival times use exponential or Poisson distributions (Bellsola Olba *et al.*, 2017; van Asperen *et al.*, 2003; Dragovic *et al.*, 2006; Shabayek and Yeung, 2002). Thus, they ignore that ships are returning with great regularity and the return times depend on the ship size. We represent ship return times in more detail as normal mixes dedicated to ship classes. This fits better the data available now thanks to AIS system.

Service time was modeled using normal (Bellsola Olba *et al.*, 2017), or Erlang (Dragovic *et al.*, 2006; Shabayek and Yeung, 2002) distributions, or was given (deterministic) (van Asperen *et al.*, 2003). In this paper, service time distributions are set independently for each ship size class and port.

In the work of Pachakis and Kiremidjian (2003) a statistical ship traffic model including vessel sizes, ship draft, number of TEUs, number of cranes and terminal revenue was developed on the basis of one port. The draft was calculated using linear regression on ship length and the number of cranes was estimated using linear regression on the number of transferred TEUs. The ship interarrival times were generated for all vessel sizes by one process with exponential distribution. We build more advanced ship size, service time and arrival time models using a much larger set of real data from multiple ports, which allows to spot port differences.

In the work of van Asperen *et al.* (2003) three arrival processes were considered: stock-controlled arrivals, equidistant arrivals, Poisson process. In the first two processes, ship arrival is planned but the actual time of arrival has a three-point distribution with probabilities: 10% of arriving [2,12] hours before the expected time, 80% of arriving within  $\pm 2$  hours around the expected time, 10% of arriving [2,12] hours after the expected time.

<sup>1</sup>Because of the rules of access we are not allowed to release the source data.

Table 1. Summary of notation.

|                  |                                                                                                                                                               |
|------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $a_i$            | number of ships of class $i$ in some port                                                                                                                     |
| $(c_{i-1}, c_i]$ | interval of ship lengths in cluster (size class) $i$                                                                                                          |
| $f_a$            | fraction of aperiodic arrivals in a certain port, or terminal                                                                                                 |
| $k$              | number of ship classes                                                                                                                                        |
| $L_j$            | length of ship $j$                                                                                                                                            |
| $\mathcal{L}_i$  | ship length model for cluster $i$ in some port                                                                                                                |
| $\nu_j$          | number of ship $j$ calls at a port in the historic data set                                                                                                   |
| $n$              | number of vessels in some port as physical objects                                                                                                            |
| $N$              | number of calls at a port (sum over vessel calls)                                                                                                             |
| $p_j$            | processing time of ship $j$ (one of possibly several instantiations in the historic data set and in STM)                                                      |
| $\mathcal{P}_i$  | model of ship processing times for cluster $i$ in some port                                                                                                   |
| $r_j$            | arrival (ready) time of ship $j$ (one of possibly several instantiations in the historic data set and in STM)                                                 |
| $\rho_j$         | return time of ship $j$ , that is duration between two arrivals of the same ship (one of possibly several instantiations in the historic data set and in STM) |
| $\mathcal{R}_i$  | model of return times for ship size class $i$                                                                                                                 |

Table 2. Example vessel data.

| vessel     | A   | B   | C   | D   | E   | F   | G   |
|------------|-----|-----|-----|-----|-----|-----|-----|
| $p_j$ [h]  | 2   | 2   | 2   | 3   | 3   | 3   | 3   |
| $L_j$ [m]  | 110 | 120 | 130 | 390 | 380 | 370 | 400 |
| Scenario 1 |     |     |     |     |     |     |     |
| $r_j$ [h]  | 0   | 3   | 6   | 0   | 3   | 6   | 9   |
| Scenario 2 |     |     |     |     |     |     |     |
| $r_j$ [h]  | 0   | 0   | 0   | 0   | 0   | 0   | 0   |
| Scenario 3 |     |     |     |     |     |     |     |
| $r_j$ [h]  | 10  | 5   | 0   | 5   | 0   | 11  | 10  |

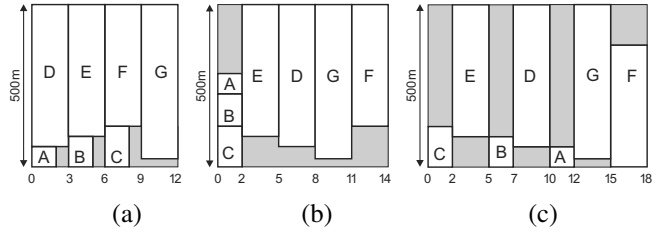


Fig. 1. Example schedules: Scenario 1 (a), Scenario 2 (b), Scenario 3 (c). Vessel and quay lengths are not proportional.

Bellsola Olba *et al.* (2018) review 18 port nautical infrastructure simulation models with the purpose of risk and capacity assessment. The extent of models and processes they represent demonstrates that modeling ship traffic is data-expensive and laborious. Hence, even partial models are helpful. They observed that uncertainty of arrival and processing times should be properly represented and recommended to base the vessel arrival times on historical data, considering the stochasticity of the process. This is what is undertaken in our work.

### 3. Features of the ship traffic model

In this section, the rationale behind the proposed ship traffic model and its relation with port logistics are presented.

We will often refer to the example of berth allocation problem (BAP), which is a central problem of port logistics. The BAP consists in assigning arriving ships to berthing positions subject to ship length and temporal constraints while optimizing some performance index (Bierwirth and Meisel, 2010; 2015).

What are the key parameters that allow us to make good decisions for this problem? Obviously, one needs to know the arrival time of each vessel. The number and sizes of the berths are a key issue, together with the length of the vessels. Finally, computing a schedule necessitates knowing the service time of each vessel. Consider the following simple example to illustrate this: A quay has one berth that is 500 m long. Vessel parameters are provided in Table 2. There are three scenarios of arrival process. Schedules for the three scenarios are

shown in Fig. 1. The first scenario is simple to manage because vessels arrive in pairs fitting quay length and it suffices to serve the ships as they arrive. In the second scenario all vessels are ready simultaneously and, to serve them, a decision of a combinatorial nature is needed. Namely, a sequence of vessel service has to be determined. The third scenario is adverse because vessels arrive in mutually incompatible pairs, which results in low quay utilization. The scenarios that appear in reality may be neither so easy nor that adverse. Moreover, just the vessel arrival process was tackled in this example. Other vessel parameters, port characteristics, or the considered optimization problem open ways for more complex relationships to emerge. Hence, there is a need for distinguishing realistic requirements on a ship traffic model. In the remaining part of this section, we discuss three main parameters defining ship traffic. The details of the building process are provided in the next section.

The arriving ship  $j$  has a certain length  $L_j$  and must be given sufficient space along a finite quay.

Large container mother ships travel long distances, e.g., on Asia–Europe or Asia–America lines, while small ships (feeders) operate on local connections. Consequently, large vessels call at the ports less frequently than the small ones. Hence, the ship length is a basic feature in an STM, and is highly related to the ship arrival time.

Temporal constraints on serving ship  $j$  are imposed by arrival time  $r_j$  and ship service time  $p_j$ , also called processing time in this paper. In the papers on the BAP as a scheduling problem, each ship arrival is often considered as an independent event. But a single vessel may call

at some port many times with some return intervals. In that case, the vessel will be called *returning*, or *periodic*. Otherwise it is called a *non-returning* ship, or a ship with *aperiodic* arrival. When a distinction between one arrival of a ship and the vessel as a physical object is important, we will use terms arrival, occurrence vs. a physical vessel/ship. Ship ready times  $r_j$  may be realizations of some return process. In this case, the physical return time of ship  $j$  will be denoted by  $\rho_j$ .

The vessel processing times  $p_j$  may be a proxy to the number of loaded and unloaded TEUs assuming that the transfer rate is roughly constant. The processing times depend on the berthing process, the quay crane schedule, yard management, intra-port transport. Let paper (Schepler et al., 2017) serve as an example of (i) port operations modeling, (ii) complexity of the relationships, (iii) problem sizes that can be currently solved. Instances with barely 48 vessels arriving over seven days, on several terminals, with intra- and inter-terminal transport were optimized. This size of instances, appropriate for tactical decisions, is far from sufficient to represent details of alternative container terminal designs at the long-term strategic level planning, with consequences spanning over several year time horizons. Thus, instead of building a planning and scheduling representation of all elements of a container terminal, it is simpler to use realistic vessel processing times (again, from the real data now available) with adequately represented dispersion. Hence, ship processing time distributions are essential components of the model.

A majority of port logistic formulations assume a deterministic approach. This means that parameters  $L_j, r_j, p_j$  are given numbers. However, it is hard to accept that  $r_j, p_j$  can be fixed in advance over time horizons typical of strategic planning. Ship arrival and processing times depend on many factors which cannot be controlled, and it is more plausible to accept that uncertainty in these values exists. For instance, more than 49% of vessels can be a day late and the average deviation from the estimated arrival time can exceed two days (Li et al., 2016). Hence, despite the advertised schedules of liner companies, arrival times of periodic ships are not deterministic and the ship traffic model should be considered as a stochastic process which samples objects  $(L_j, r_j, p_j)$ .

Let us return to the interrelations between  $L_j, r_j, p_j$  parameters. The service given to a big ship has a bigger value for the terminal operator than the service given to a smaller ship. Hence, there are stronger financial incentives to stick to the schedules planned for the big ships. This can be reflected in different dispersions of ship processing and arrival times.

As already said, ship return and processing times depend on the ship sizes. In order to deal with such dissimilarity, different ship length *classes*, or equivalently

called *clusters*, will be considered. Each cluster should have its own distribution of  $p_j$  and  $r_j$  values.

Finally, let us note that in the intended STM, we want to identify and generalize traffic patterns important for long-time strategic planning. It implies that short-term weekly and daily data patterns are less interesting as typical of contemporary operational optimization. We will return to this issue in Section 4.5.2.

## 4. Building the ship traffic model

The ship traffic model is defined by the following elements: (i) ship class definitions, (ii) a model  $\mathcal{L}_i$  of ship lengths for each class  $i$ , (iii) a model of processing times  $\mathcal{P}_i$  for each class  $i$ , (iv) ready time models: return time models  $\mathcal{R}_i$  for each class  $i$  of returning ships, and a ready time model for non-returning ships. Thus, the method to build a ship traffic model on the basis of the actual data consists of three steps: (i) partitioning of the ships into size classes, (ii) setting processing time distributions for each size class, (iii) building return, or ready time models. In the next subsection, the historical data set is described. Then, the STM parts are derived. After proposing components of the STM, a discussion of the model limitations and alternative approaches will be conducted.

**4.1. Data set.** The ship traffic model was built from the historical automatic identification system (AIS) data gathered in eight ports in the world in 2016. The data set comprises a range of port sizes, from small traffic ports (Gdańsk) to the largest in the world (Singapore). Basic data set information is collected in Table 3. It comprises the number of calls at the ports, the number of physical vessels, returning vessels, the number of physical vessel lengths given within 1 m resolution, and fraction  $f_a$  of aperiodic ships in the total number of arrivals. For example, in Gdańsk only 24 calls out of 471 were single visits. The fraction of non-returning calls ranges from 1.21% (Singapore), to 6.25% (Le Havre). Table 3 provides the first qualitative conclusion for port design and port logistic algorithm testing: a large majority of the vessels are returning. Consequently, the cardinality of physical ship set is rather limited.

**4.2. Ship size clustering.** The way of exploiting a ship and handling it at the terminal depends on its size  $L_j$ . Despite attempts to limit the number of ship classes (see Panamax, NewPanamax and so on), there is a large variety of vessel sizes (cf. the number of unique  $L_j$ s in Table 3 and examples in Fig. 2), but for practical reasons it is more convenient to use a few vessel size classes. Indeed, the ship traffic model dedicated to a set of similar ships is simpler, easier to develop and more accurate than a general model encompassing all possible

Table 3. Basic data set information.

| port                 | Gdańsk  | Long Beach | Los Angeles | Le Havre  |
|----------------------|---------|------------|-------------|-----------|
| number of calls $N$  | 471     | 995        | 1308        | 2271      |
| physical ships $n$   | 81      | 242        | 309         | 559       |
| returning ships      | 57      | 188        | 238         | 417       |
| No.of unique $L_j$ s | 31      | 65         | 61          | 105       |
| $f_a[\%]$            | 5.10    | 5.43       | 5.43        | 6.25      |
| port                 | Hamburg | Rotterdam  | Shanghai    | Singapore |
| number of calls $N$  | 3294    | 3998       | 11606       | 18494     |
| physical ships $n$   | 586     | 311        | 1233        | 1857      |
| returning ships      | 466     | 240        | 1038        | 1634      |
| No.of unique $L_j$ s | 101     | 80         | 148         | 169       |
| $f_a[\%]$            | 3.64    | 1.78       | 1.68        | 1.21      |

ship lengths. We will show in Section 4.4 that a single size-general model for ship processing times  $p_j$  appeared unsatisfactory. Hence, vessels calling at a particular port were clustered according to their lengths.

The definition of a ship class (cluster)  $i$  consists in the ship size range. The number  $a_i$  of arrivals in the class can then be calculated. Let  $(c_{i-1}, c_i]$  be the range of vessel sizes for the  $i$ -th cluster, where  $c_0 = 0$  and  $c_k = \max_j \{L_j\}$  for the last cluster  $k$ . The quality of vessel  $j$  fit in cluster  $i$  is measured by the distance

$$d(j, i) = \begin{cases} \infty & \text{for } L_j > c_i, \\ c_i - L_j & \text{for } L_j \leq c_i. \end{cases} \quad (1)$$

Let  $\mathbb{1}(j, i) = 1$  if  $d(j, i)$  is minimum for ship  $j$  and cluster  $i$ , and 0 otherwise. The quality of clustering is measured by the sum of all ship distances weighted by the number of calls  $\nu_j$  of ship  $j$  at the port considered:

$$Q(c_1, c_2, \dots, c_k) = \sum_{j=1}^n \nu_j \mathbb{1}(j, i) d(j, i). \quad (2)$$

While setting the clusters, their number  $k$  and range ends  $c_i$  are the decision variables. Note that since the quality of clustering is weighted by  $\nu_j$ , some groups of ships with many returning calls may be split into clusters of a narrower range  $(c_{i-1}, c_i]$ , while some less frequent groups of ships can be merged together. The  $c_i$  values were calculated by the generalized reduced gradient method (Lasdon *et al.*, 1974) for  $k = 5, 6, 7$ . The values of the clustering quality as measured by (2), normalized to the best result, are shown in Table 4. Since the quality of clustering does not improve much by raising  $k$  from 6 to 7, and since using many ship size classes is unwieldy, we decided to limit the number of ship sizes clusters to  $k = 7$ .

The ends  $c_i$  of the size intervals for the ports considered are given in Table 5. Additionally, the total

number  $a_i$  of calls for cluster  $i$  and the number of aperiodic arrivals are given in brackets. In the following we will refer to the clusters using a short-hand notation consisting of the initials of the port name and the cluster number starting from 1 for the shortest ships. For example, RT3 is the third cluster for Rotterdam, with 532 calls including 2 aperiodic arrivals (cf. Table 5).

**Discussion.** It can be seen in Table 5 that size classes are distinctive for each port. However, some similar clusters can be identified among large ships: {GD7, LB7, LH7, HB7}, {GD6, LH6}, {LB5, LA6}, {SH6, SI6}, {LB6, HB6}. The clustering obtained has its peculiarities. For example, there is only one ship in LB7 and she is aperiodic. This cluster will require a special handling: the ship data can be used directly as a representative for the LB7 cluster.

We built different size ranges  $(c_{i-1}, c_i]$  for each port. Alternatively, the same ranges can be defined for all ports in the world. We will present ship sizes from this perspective in Section 6 (cf. Fig. 10). An advantage of this alternative is that it makes port comparisons easier. A drawback is that some of such classes may be empty in some ports and the number of classes must be bigger to sufficiently discern ship size differences between ports, which complicates the STM. We decided to use size ranges specific to each port to get smaller numbers of classes and simpler STMs.

**4.3. Ship length model.** The easiest way of representing ship lengths in a certain class is to unify the whole cluster and to use the upper end of the cluster range. That is, model  $\mathcal{L}_i$  for a cluster with range  $(c_{i-1}, c_i]$  is  $c_i$ . Let us call this representation the longest-ship model.

**Discussion.** A more precise analysis of ship lengths in clusters reveals data and clustering peculiarities. In

Table 4. Normalized clustering quality scores.

| port    | Gdańsk  | Long Beach | Los Angeles | Le Havre  |
|---------|---------|------------|-------------|-----------|
| $k = 5$ | 2.10    | 1.88       | 1.57        | 1.91      |
| $k = 6$ | 1.45    | 1.33       | 1.15        | 1.46      |
| $k = 7$ | 1       | 1          | 1           | 1         |
| port    | Hamburg | Rotterdam  | Shanghai    | Singapore |
| $k = 5$ | 1.91    | 1.50       | 1.42        | 1.52      |
| $k = 6$ | 1.09    | 1.25       | 1.06        | 1.31      |
| $k = 7$ | 1       | 1          | 1           | 1         |

Table 5. Ship size cluster interval end  $c_i$  in meters, the total number  $a_i$  of calls in the clusters and number of aperiodic calls.

| port  | Gdańsk       | Long Beach    | Los Angeles   | Le Havre      |
|-------|--------------|---------------|---------------|---------------|
| $c_1$ | 137 (115,5)  | 188 (89,1)    | 224 (167,9)   | 140 (216,6)   |
| $c_2$ | 151 (103,2)  | 232 (137,2)   | 261 (132,4)   | 210 (337,25)  |
| $c_3$ | 183 (121,1)  | 273 (118,11)  | 279 (180,5)   | 245 (333,7)   |
| $c_4$ | 210 (11,3)   | 302 (203,14)  | 295 (252,12)  | 278 (400,19)  |
| $c_5$ | 300 (17,10)  | 338 (254,15)  | 305 (181,12)  | 300 (355,29)  |
| $c_6$ | 368 (50,2)   | 368 (193,10)  | 335 (294,13)  | 368 (528,49)  |
| $c_7$ | 399 (54,1)   | 399 (1,1)     | 399 (102,16)  | 399 (104,7)   |
| port  | Hamburg      | Rotterdam     | Shanghai      | Singapore     |
| $c_1$ | 141 (905,8)  | 102 (153,2)   | 101 (642,2)   | 172 (4552,21) |
| $c_2$ | 170 (956,6)  | 141 (2395,13) | 148 (4144,11) | 198 (3206,21) |
| $c_3$ | 213 (298,5)  | 152 (532,2)   | 183 (2263,20) | 225 (2231,21) |
| $c_4$ | 279 (302,23) | 170 (444,4)   | 237 (1584,30) | 262 (2182,47) |
| $c_5$ | 338 (339,45) | 223 (240,26)  | 297 (1800,65) | 302 (3188,67) |
| $c_6$ | 369 (422,27) | 273 (154,13)  | 348 (1022,53) | 345 (1635,28) |
| $c_7$ | 400 (72,6)   | 305 (80,11)   | 367 (151,14)  | 400 (1500,18) |

Fig. 2 empirical cumulative distribution functions (CDF) of the ship lengths in example clusters are shown against theoretical CDFs of selected continuous distributions fitting ship lengths best. Parameters of the probability distributions chosen by the `fitdistrplus` method of R language are available from Wawrzyniak et al. (2021).

Even without a close look at these distributions, Fig. 2 shows that there are clusters which follow very well the above longest-ship approach. For example, the ship lengths in cluster LH7 come from only five lengths in range [395, 399] m (Fig. 2(a)) and rounding them up to 399 m is not a substantial loss of precision. There are also clusters which would be better split into a few sub-clusters. For instance, visually LB6 (Fig. 2(b)) could be much better divided into clusters of lengths  $\approx 350$  m and  $\approx 365$  m. Similar results were obtained for LA7 (not shown here). Thus, a better ship length model would be possible if more size classes were used for particular ports. However, this leads to a dilemma whether to emulate particular ports in more detailed way (at the cost of more complex model), or on the contrary, generalize the results.

Since our goals are in generalizing patterns in ship traffic, we do not follow the first path. Finally, there are clusters for which ship lengths can be quite well approximated by continuous distributions. This is the case for SI1 and SH2 (Figs. 2(c) and (d)). However, the reader should be aware that for some clusters (e.g., as narrow as LH7) searching for a continuous ship length distribution is unfounded.

**4.4. Processing time model.** Our first attempt to develop a vessel processing time model consisted in calculating a linear regression function of the processing time in the ship length:  $p_j = a_1 L_j + a_2$ , for all the ships of the port considered. In Fig. 3 examples of ship processing times  $p_j$  and processing time per unit of ship length  $p_j/L_j$  are shown. In all figures, the linear regression line, its parameters, and the coefficients of determination  $R^2$  are also given. It can be observed that the distributions of  $p_j$  and  $p_j/L_j$  depend on the ship size. Usually, longer ships have larger  $p_j$ , but a negative correlation between  $p_j/L_j$  and  $L_j$  can be observed.

However, it is obvious from Fig. 3 that linear

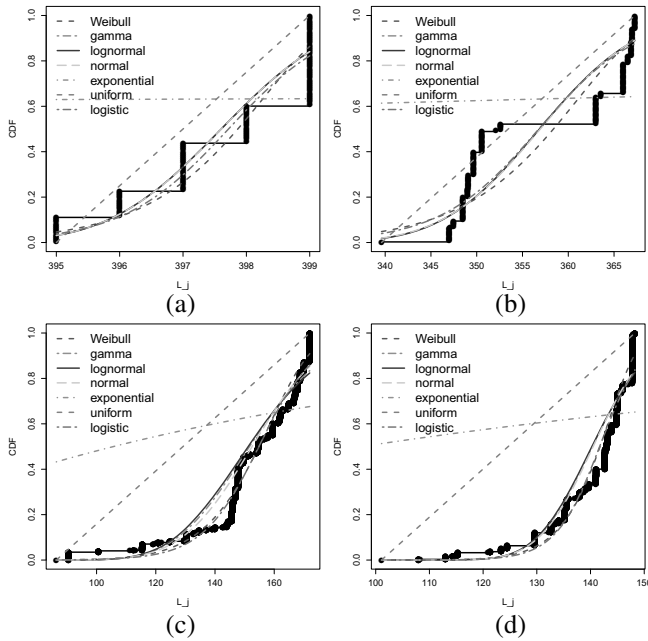


Fig. 2. Empirical (bold black) vs. theoretical ship length distributions in example clusters: LH7 (a), LB6 (b), SI1 (c), SH2 (d).

regression is not a good predictor of  $p_j$ . There is a great deal of dispersion both in  $p_j$  and in  $p_j/L_j$ . This is reflected in very low coefficients of determination  $R^2$ . Thus, linear regression is not suitable to reproduce accurately the processing times. It can be seen that  $p_j$  values align vertically in a way corresponding to ship size classes or particular ships. Hence, it is advisable to analyze ship processing times for distinct ship size classes, rather than using a single model for all possible lengths. In the following, we decided to consider  $p_j/L_j$ , rather than  $p_j$  distributions because the former are more compact and their ranges are more similar between the ports. Furthermore,  $p_j/L_j$  can be considered a better score for the logistic technologies used in the port than  $p_j$ . Correlations between  $L_j$  and  $p_j/L_j$  can indicate the economy of scale effect for the port considered. This will be discussed in Section 6.

Eight common parametric probability distributions (beta, exponential, gamma, normal, lognormal, logistic, uniform, Weibull) were fit to the  $p_j/L_j$  values for ship size clusters of the ports considered, by using the `fitdistrplus` package of the R language (Delignette-Muller and Dutang, 2015). Distribution parameters were calculated using maximum likelihood estimation. Examples of `fitdistrplus` package visual output are shown in Fig. 4. Visual data analyses are often unrestrictive, while numeric goodness-of-fit statistics often disagree in their recommendations. Moreover, we operate on a data set representing a real ship stream which is rich in information and peculiarities.

To make the STM building process deterministic and reproducible, we used the Anderson–Darling goodness-of-fit score provided by `gofstat` from the R programming language to choose the best fitting probability distribution for each port and each ship size class. The distributions selected are summarized in Table 6. The parameters of the best fitting distributions are provided by Wawrzyniak *et al.* (2021).

**Discussion.** For each distribution, the number of size clusters that are the best fit (wins), the number of clusters that are the worst fit (worst), and the best fit clusters, are given in Table 6. The best fit clusters are the ones for which the distribution considered had the best fit according to the Anderson–Darling goodness-of-fit score. The number of worst fits is the number of cases for which the distribution considered provided the worst fit. The maximum number of possible wins is 55, i.e., the number of all ship clusters in all ports but LB7 which has exactly one aperiodic ship. Let us note that uniform distribution was excluded from Table 6 because it is always the worst choice. Exponential and normal distributions are the second and third worst after the uniform. Uniform, exponential, and normal distributions very often used in simulation and analytical modeling perform bad here, which is an interesting qualitative observation for modeling scheduling problems or testing algorithm performance. Although the lognormal and the logistic probability distributions fit best the  $p_j/L_j$  distributions in many cases, none of the tested distributions is predominantly the best. This lack of regularity between ports for  $p_j/L_j$  winning distributions demonstrates that each port is unique when the processing time of the ships is considered.

**4.5. Ready time model.** Since most of the arrivals are periodic (cf. Table 3), we start with the returning ships. The non-returning ships will be dealt with in the following subsection.

**4.5.1. Returning ships.** In Fig. 5 examples of ship return times  $\rho_j$  in days are shown for Le Havre, Los Angeles, and Singapore. In Fig. 5(a) return times  $\rho_j$  are presented vs. ship lengths  $L_j$ , and in Fig. 5(b) histograms of the return times are given. The most frequent ship return intervals can be identified as week multiplicities in Fig. 5(b), which is typical of shipping network design practices. It is expected that return times depend on the ship size class. Usually, the longer the ship, the longer the return times (Fig. 5(a)). This is confirmed by positive correlation between  $L_j$  and  $\rho_j$  values. The small values of  $R^2$ , and the observed large variation in the return

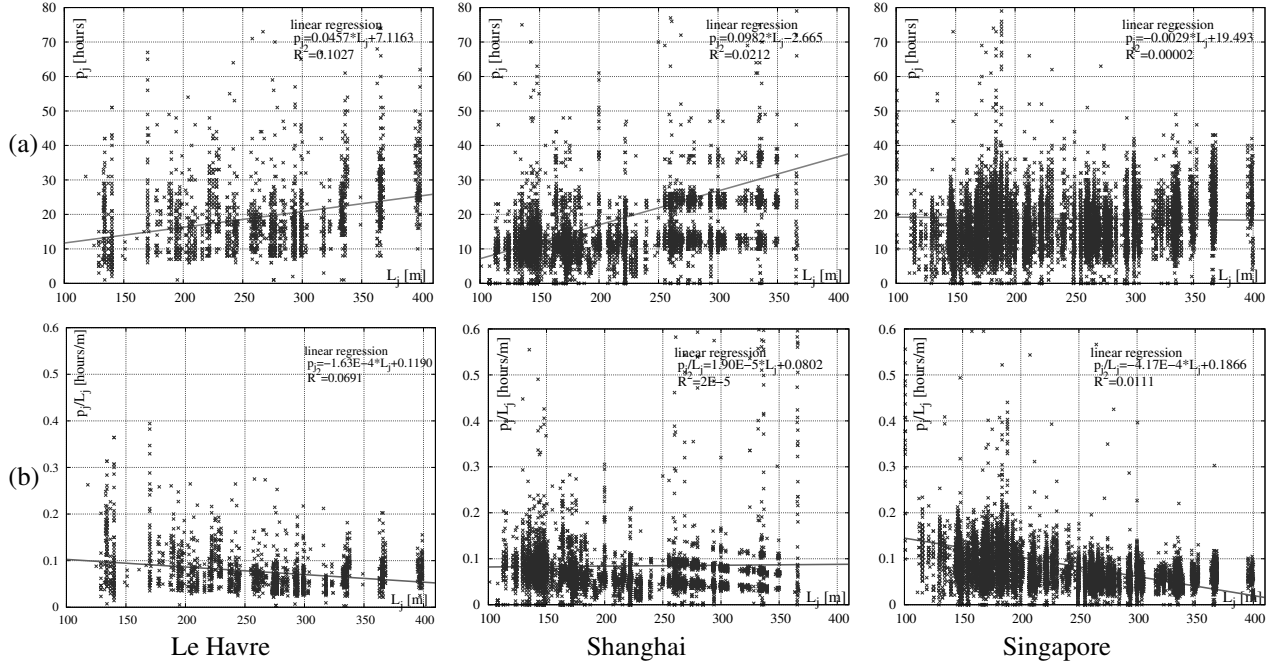


Fig. 3. Examples of relationships:  $p_j$  vs.  $L_j$  (restricted to  $p_j \leq 80$ ) (a),  $p_j/L_j$  vs.  $L_j$  (restricted to  $p_j/L_j \leq 0.6$ ) (b). Le Havre (left), Shanghai (center), Singapore (right). Linear regression fit and coefficients of determination are as follows: Le Havre:  $p_j = 0.0457L_j + 7.1163$ ,  $R^2 = 0.1027$  (a),  $p_j/L_j = -0.000163L_j + 0.1190$ ,  $R^2 = 0.0691$  (b); Shanghai:  $p_j = 0.0982L_j - 2.655$ ,  $R^2 = 0.0212$  (a),  $p_j/L_j = 0.000019L_j + 0.0802$ ,  $R^2 = 2E-5$  (b); Singapore:  $p_j = -0.0029L_j + 19.493$ ,  $R^2 = 2E-5$  (a),  $p_j/L_j = -0.000417L_j + 0.1866$ ,  $R^2 = 0.0111$  (b) ( $p_j$  in hours,  $L_j$  in meters).

times, suggest that other parameters are involved, e.g., shipping network timetables, weather conditions or ship processing time variability. Hence, a model assuming some fixed return period would not fit well the observed large scattering of the return times.

In order to construct a model  $\mathcal{R}_i$  of returns in cluster  $i$ , we applied method `normalmixEM` from R package `mixtools` (Benaglia et al., 2009), to fit a mixture of normal distributions into return times of each ship size class for each port. Assume that a set of ship return observations  $\bar{p}_i = [\rho_1, \dots, \rho_{a_i}]$  for cluster  $i$  is given, where  $a_i$  is the number of returns in cluster  $i$ . The method `normalmixEM` fits a probability density function  $g$  which for some return value  $\rho_j$  can be written as

$$g(\rho_j|\bar{\theta}) = \sum_{h=1}^{\ell} \lambda_h \phi_h(\rho_j|\mu_h, \sigma_h^2), \quad (3)$$

where  $\ell$  is the number of components in the mixture,  $\phi_h$  is the normal probability density function with mean  $\mu_h$  and variance  $\sigma_h^2$ , values  $\lambda_h$  are mixing proportions which are positive with a sum equal to 1,  $\bar{\theta}$  is a density mixture parameter vector comprising  $\ell$  triplets  $(\lambda_h, \mu_h, \sigma_h^2)$ . Mixture parameters  $\bar{\theta}^*$  are chosen to maximize the fitting quality which is the logarithm of

likelihood (*loglik* for short) of the data obtained:

$$\bar{\theta}^* = \arg \max_{\bar{\theta}} \sum_{j=1}^{a_i} \log g(\rho_j|\bar{\theta}). \quad (4)$$

The obtained mixture parameter vector  $\bar{\theta}^*$  is a local, rather than a guaranteed global optimum. The run parameters of `normalmixEM` were set to default values, except for `maxit=50000`, `maxrestarts=200` which were chosen experimentally to obtain results for the widest set of ship size classes and component numbers  $\ell$ . Construction of mixtures with  $\ell = 2, \dots, 20$  was attempted. The number of components  $\ell$  which provided maximum loglik in the verified range of values was chosen as the best mixture. Examples of visual results for fitting return times with mixtures (3) are shown in Figs. 6 and 7. The summary of return time models  $\mathcal{R}_i$  for all clusters  $i$  can be found in the work of Wawrzyniak et al. (2021).

**Discussion.** It is known that finding a matching mixture can be challenging (Benaglia et al., 2009). Feasibility of `normalmixEM` depends on, e.g., the size  $a_j$  of the data sample, the number of mixture components  $\ell$  and the actual dispersion of the data. And indeed, mixtures could not be obtained for all ports, clusters and  $\ell$  values. In particular for Gdańsk the range of feasible mixture components was the narrowest, often limited to just

Table 6.  $p_j/L_j$  distributions selected for the ship size classes and ports.

| Distribution | No. of wins | No. of worst | Selected (winning) clusters                                                             |
|--------------|-------------|--------------|-----------------------------------------------------------------------------------------|
| beta         | 5           | 0            | GD4,RT5,LA4,SI5, SI6                                                                    |
| exponential  | 1           | 49           | SH1                                                                                     |
| gamma        | 10          | 0            | GD3, LH6, LH7, RT1, RT3, RT4, LA6, SI1, SI3, SI4                                        |
| normal       | 2           | 4            | HB2, LA2                                                                                |
| log-normal   | 18          | 2            | GD1, GD5, GD6,LH1, LH2, LH3, LH5, HB3, HB4, HB6, RT2, RT6, LA1, LB4, SH5, SH6, SH7, SI2 |
| logistic     | 14          | 0            | LH4, HB5, RT7, LA3, LA5, LA7, LB1, LB3, LB5, LB6,SH2, SH3, SH4, SI7                     |
| Weibull      | 5           | 0            | GD2, GD7, HB1, HB7, LB2                                                                 |

$\ell = 2, 3$ . Moreover, there are only 4 different return periods in GD5. In most of the cases, the fitness quality (loglik) improves with the increasing number of mixture components  $\ell$ . A large number of components in (3) is impractical; therefore, it is an attractive idea to limit the number of used mixture components. However, it is hard to choose a threshold of  $\ell$  in an indisputable way. For example, in Fig. 7(a) changes of loglik relative to the best obtained value (at  $\ell = 20$ ) and density functions for  $\ell = 11, 16, 20$ , are shown for SI7. Though all shown density functions (Figs. 7(b)–(d)) at least visually cover the data well, the range of loglik minimization is below 10%, but the value of loglik decreases slowly with  $\ell$  (Fig. 7(a)) and it is hard to point a single incontrovertible value of  $\ell$  at which the process of increasing the number of components could be stopped. For reproducibility we stuck to the choice of  $\ell \leq 20$  for which the loglik was the smallest.

On the basis of the identified return intervals, qualitative observations can be made for terminal planning and scheduling. Our study confirms that large ships call ports according to schedules with week multiplicities, but shorter schedules also exist. It can be seen in Fig. 5(b), that 6-, 7-, 8-, 11-week return times are quite common. Thus, one-week planning horizons (common in practice) are adaptations that hardly represent a more complex process. In order to grasp interactions between different return periods, the least common multiple of the periods should be considered as a planning horizon. This may easily span more than a year.

We provided a probabilistic model of return times  $\mathcal{R}_i$  for cluster  $i$  of a given port. It remains to define the first arrival. Several approaches are possible: (i) draw

at random from normal distribution with mean  $\mu_k$  and variance  $\sigma_k^2$  where  $k = \arg \max_{\ell} \{\lambda_{\ell}\}$ ; (ii) draw returns many times from model (3) as if returning over many years to accumulate dispersion and use the offset from the beginning of the current year; (iii) draw a random day of the week (DoW) and then draw a random hour of the day (HoD) using arrival distributions from Wawrzyniak *et al.* (2021); (iv) combine methods by first using (i) and then moving the first arrival to the nearest day and hour generated by (iii).

**4.5.2. Non-returning ships.** Though the majority of the calls at ports are returning, aperiodic arrivals are not exceptions and there are over 50 ships in most of the analyzed ports, which is more than one arrival a week. Hence, we decided to include this group of ships in the traffic model. The number of aperiodic ships in each size cluster was given in Table 5. Aperiodic arrivals are not concentrated in a restricted subset of clusters. Hence, ship classes and lengths can be modeled in the same way for all ships, both returning and non-returning. As the processing time model  $\mathcal{P}_i$  for non-returning ships, we propose to use the model of the corresponding size class (Section 4.4). Note that even for the clusters with the biggest number of aperiodic arrivals (see Table 5, SH6: 53, SH5: 65, SI5: 67) on the average there is slightly more than one aperiodic ship per week in a size class. The median number of aperiodic ships per week in one size class calculated over all ports and clusters is 0.211 which is roughly equivalent to one aperiodic arrival per cluster over 33 days ( $\approx 4.7$  weeks). Hence, aperiodic arrivals are not common enough to build *separate* and trustworthy statistical models of ready times for each size class with high time resolution.

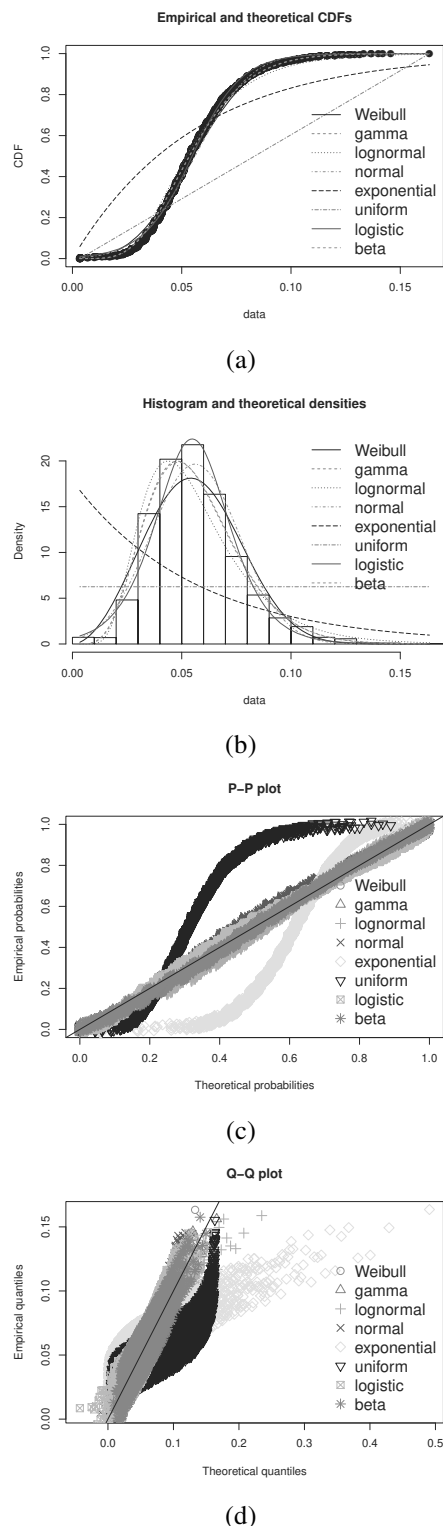


Fig. 4. Example of visual output from `fitdistrplus` for SI5: cumulative distribution functions, where black points are the actual observations (a), probability density functions (b), PP plots (c), QQ plots (d).

Therefore, to model the non-returning ship, arrival process with a similar resolution as for the returning ships we will apply a superposition of three distributions: for the week of a year, for the day of a week *DoW*, and for the hour of a day *HoD*. This was done for whole ports, without distinguishing size classes.

**Weeks of the year.** It appeared that the number of weekly aperiodic arrivals changes over the year and in most of the cases the number of aperiodic arrivals increases at the turns of the year (see Fig. 8(a)). Thus, aperiodic arrivals show seasonal fluctuations. In order to represent the variability over the weeks of the year, we decided to build an equivalent of a density function by fitting a polynomial  $f(w)$  of the week number  $w$  into the weekly frequencies.

A careful study led to use of fourth-degree polynomials, which is the lowest even degree allowing a smooth transition between consecutive years. The visual output of the fitting can be verified in Fig. 8(a). Then specific 4th degree polynomials were obtained by fitting the data for each port. Examples of visual output can be seen in Figs. 8(b)–(d). The coefficients of the polynomials fitting the relative frequencies of aperiodic arrivals are given by Wawrzyniak et al. (2021).

**Days of a week and hours of a day.** Examples of the distributions of aperiodic arrivals over days of a week for Le Havre and Rotterdam are shown in Fig. 9(a). Along the vertical axis fractions of the total number of aperiodic arrivals are shown. In Fig. 9(b) the coefficient of variation for aperiodic arrivals over days of a week is shown vs. the number of aperiodic arrivals in the port during the year. Results for Le Havre are shown in Fig. 9(a) because they present an appealing A-shaped pattern with the top arrivals on Wednesdays. In Rotterdam, Fridays were days with a  $2.15 \times \sigma$  departure from the average, where  $\sigma$  is the standard deviation of the daily fractions in the port. This was the biggest departure from the average in all studied ports. It may be the case that these two ports exhibited some tendency toward processing non-returning ships on these two days. However, it may be also a human perception artifact. As many as 64% of aperiodic arrivals of all ports on days of a week fit in the  $1 \times \sigma$  range around the average. The exceptional “Rotterdam Friday” may emerge randomly with a probability of 0.019 if the dispersions of aperiodic daily arrivals is normally distributed. Moreover, it can be seen in Fig. 9(b) that ports with a larger number of aperiodic ships have smaller dispersion of arrivals between days of a week. Hence, there are good reasons to think that the number of aperiodic arrivals changes randomly over the days of a week. Therefore, we applied the `fitdistrplus` method from the R programming environment to find the distribution fitting best the variability of fractions of

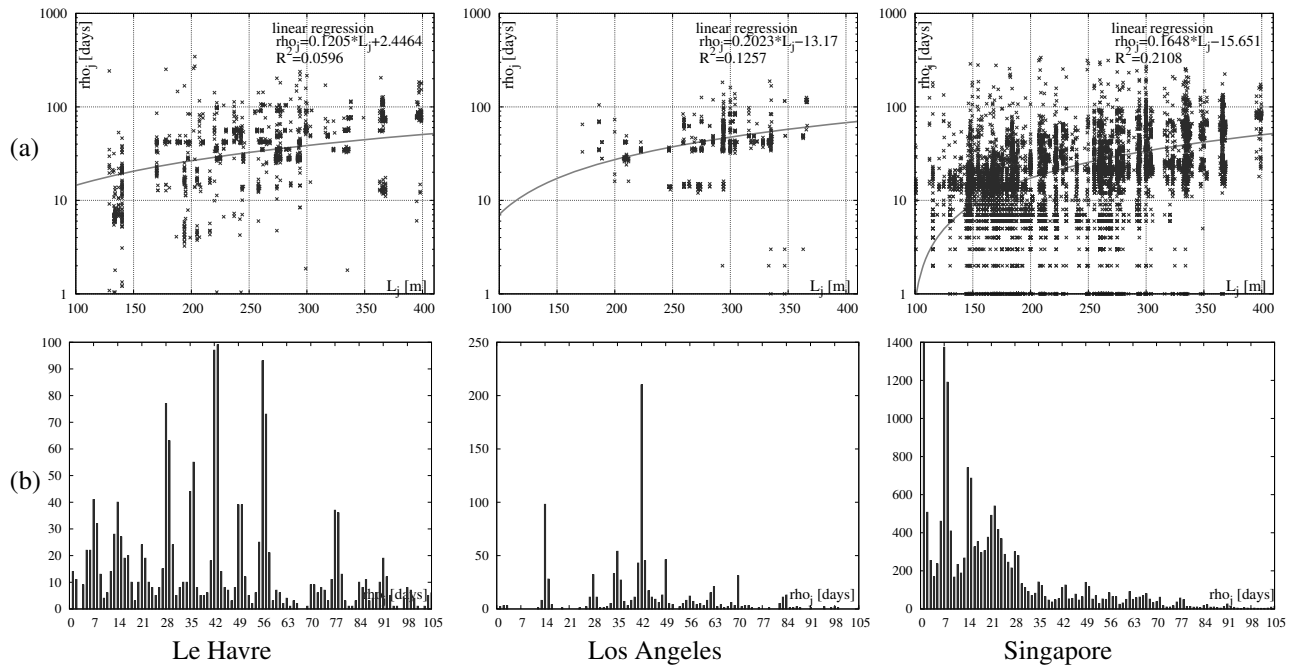


Fig. 5. Examples of return intervals for Le Havre (left), Los Angeles (center), Singapore (right): return intervals vs.  $L_j$ ; linear regression and coefficients of determination are shown in the pictures (the vertical axis is logarithmic and hence the linear function is not a straight line here) (a), histogram of return intervals (b). Here  $\rho_j$  linear regression equations and coefficients of determination are as follows: Le Havre:  $\rho_j = 0.1205L_j + 2.4464$ ,  $R^2 = 0.0596$ ; Los Angeles:  $\rho_j = 0.2023L_j - 13.17$ ,  $R^2 = 0.1257$ ; Singapore:  $\rho_j = 0.1648L_j - 15.651$ ,  $R^2 = 0.2108$  ( $\rho_j$  in days,  $L_j$  in meters).

aperiodic arrivals over days of a week. For example, the best fit of aperiodic arrivals over days of a week aggregated over all ports is the logistic distribution. Similar analysis was conducted for aperiodic arrivals over hours of a day. Results are collected by Wawrzyniak *et al.* (2021).

**Discussion.** We proposed a 4th degree polynomial of the week number in a year as ready time model. This function can be used as an empirically-built probability density function of aperiodic arrivals in a given week. By applying a continuous function we attempted to extract a smoothed pattern in the central tendencies of weekly arrivals. However, an aperiodic arrival distribution can be defined in alternative ways, e.g., as a list of probabilities for particular weeks averaged over ports. An advantage of such a representation is simplicity. A disadvantage is the lack of smoothing and a generalizing effect of a continuous function. The lack of data for multiple years impedes constructing more reliable models, both continuous and time-discretized.

The 4th degree of the fitting polynomial is a minimum which allows smooth transition between years, which seems a practical requirement. But for reasons of computational simplicity, lower degree polynomials may also be acceptable. A visual comparison of the 2nd and 3rd degree polynomials in Fig. 8(a) suggests that the

advantage of fitting as large as a 4th degree polynomial is minor.

Due to the shortage of data, the quality of models for aperiodic arrivals over days of a week (DoW) and hours of a day (HoD) is low. Hence, for the reason of confidence in the models, it seems more legitimate to use analogous models for all ships, i.e., both periodic and aperiodic. It may be also considered advantageous because with respect to assigning a service time in a day of the week and an hour of the day, aperiodic ships are handled in the same way as periodic ships. HoD and DoW distributions both for periodic and for aperiodic arrivals are available from Wawrzyniak *et al.* (2021).

Let us conclude about arrival time modeling. In the ship arrival time models, various decision levels intersect. On the one hand, it was possible to build a continuous model for periodic arrivals. This model incorporates and generalizes traffic patterns resulting from, e.g., a global network of main shipping lines or delays due to weather conditions. On the other hand, due to the lack of data and a need for representing aperiodic arrivals with a resolution comparable to the periodic arrivals, we resorted to the use of DoW and HoD models. Yet, the weekly and daily operations of ports are subject of online optimization these days, and hence, the DoW and HoD models may change and become deterministic by nature.

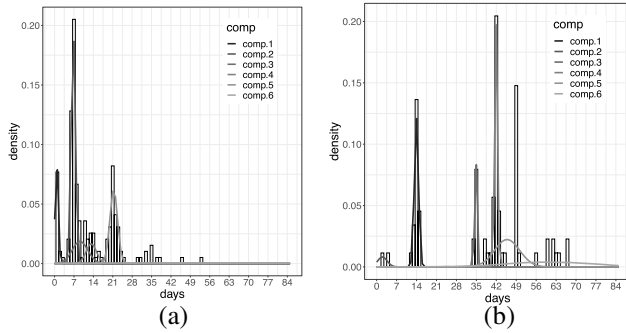


Fig. 6. Examples of return times mixtures for LH1 (a), LB3 (b) (horizontal axes in days).

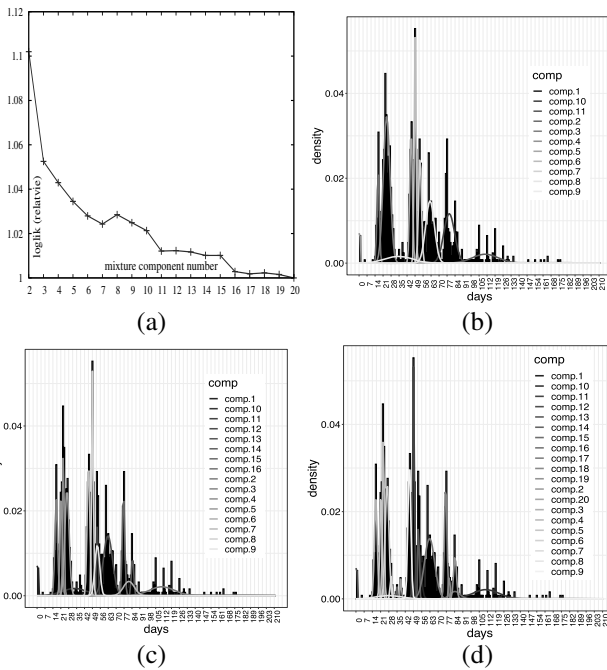


Fig. 7. Examples of loglik changes with mixture component number  $\ell$  in SI7 loglik vs.  $\ell$  (a); density functions for  $\ell = 11$  (b),  $\ell = 16$  (c),  $\ell = 20$  (d).

## 5. Applying the model

In this section we describe how to apply our STM to generate instances of the ship traffic with the above features. The consecutive steps are represented as Algorithm 1.

Algorithm 1 accepts the ship traffic model  $STM$ , the number of calls  $N$ , and the interval  $H$  of arrivals as input parameters. Particular ports are emulated by the input  $STM$ . The intensity of ship traffic, measured by the number of arrivals in some time interval, is regulated by the number of calls  $N$  and the interval  $H$  of arrivals. The clusters (size classes) are generated in loop 2–14. The number of arrivals  $a_{cl}$  for cluster  $cl$  is calculated using the numbers of arrivals  $a_{cl}^{STM}$  in the STM built on historical data. Arrivals are generated in

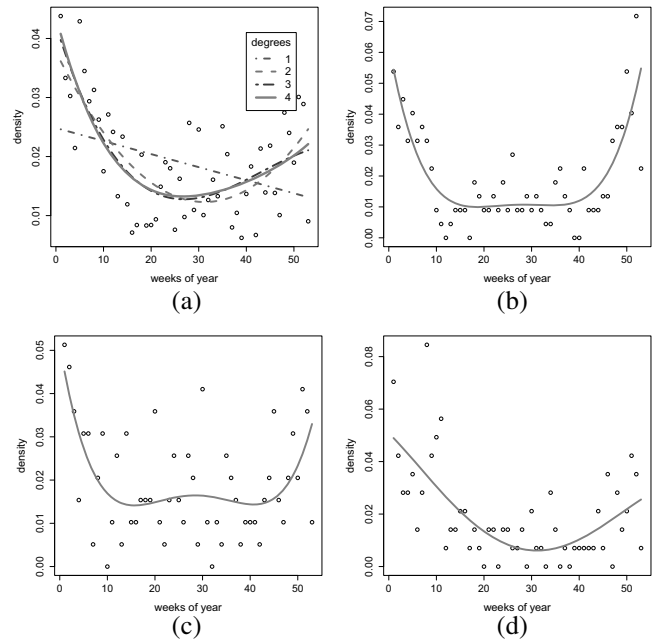


Fig. 8. Examples of fitting polynomials for aperiodic arrivals over year's weeks: fitting various degree polynomials into the normalized sum of frequencies in all ports (a); the weekly frequencies and the best-fit 4th-degree polynomial for Singapore (b), Shanghai (c), Le Havre (d).

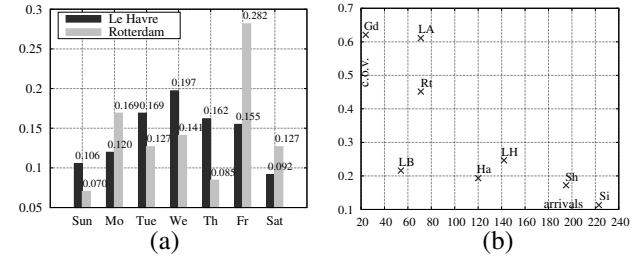


Fig. 9. Examples of aperiodic arrivals over days of a week: fractions of arrivals for Le Havre and Rotterdam (a), coefficient of variation vs. the number of aperiodic arrivals per year for different ports (b).

loop 4–14 and  $j$  is the total arrival counter. The size of the arriving ship  $L_j$  is chosen in Step 5. The modeler has to make a design decision here on the generation method ship length  $L_j$ , according to the options outlined in Section 4.3. The type of arrival is chosen in Step 6: it is an aperiodic arrival with probability  $f_a$ , otherwise it is a returning ship.  $\text{Rnd}(1)$  is a pseudo-random number generator providing numbers in range  $[0, 1]$  with uniform distribution. Function  $\text{DrawAperiodic}$  generates data for an aperiodic arrival in cluster  $cl$  of a ship with length  $L_j$ . Function  $\text{DrawReturning}$  generates at most  $a_{cl}$  arrivals of the returning ship in cluster  $cl$  of length  $L_j$  as sequences of arrivals. The two functions return the number of generated arrivals stored in  $\text{ArrNum}$ . The

**Algorithm 1.** Instance generator( $STM, N, H$ ).

---

```

1:  $j \leftarrow 1$ ;
2: for  $cl$  in 1 to 7 do
3:    $a_{cl} \leftarrow N \times a_{cl}^{STM} / \sum_{k=1}^7 a_k^{STM}$ ;
4:   while  $a_{cl} > 0$  do
5:     choose length  $L_j$  of the ship in call  $j$  as
       described in Section 4.3;
6:     if  $\text{Rnd}(1) \leq f_a$  then
7:        $\text{ArrNum} \leftarrow \text{DrawAperiodic}(j, cl, L_j)$ 
8:     else
9:        $\text{ArrNum} \leftarrow \text{DrawReturning}(j, cl, L_j, a_{cl})$ 
10:    end if
11:     $j \leftarrow j + \text{ArrNum}$ ;
12:     $a_{cl} \leftarrow a_{cl} - \text{ArrNum}$ ;
13:  end while
14: end for

```

---

**Algorithm 2.** DrawAperiodic( $j, c, L_j$ ).

---

```

1: generate number  $x$  according to the continuous
   distribution of ratios  $p_c/L_c$  for cluster  $c$  (Section 4.4),
   set processing time  $p_j = x \times L_j$ ;
2: choose arrival date  $r_j$  according to the model in
   Section 4.5.2;
3: record  $(j, L_j, p_j, r_j)$ 
4: return 1;

```

---

global arrival counter  $j$  is increased and the remaining number of arrivals  $a_{cl}$  is decreased in Steps 11 and 12, respectively.

Function DrawAperiodic is shown as Algorithm 2. In Steps 1 and 2, processing time and arrival date are generated. In Step 3, a tuple of values defining the  $j$ -th arrival is recorded.

Function DrawReturning is defined as Algorithm 3. In DrawReturning a sequence of the same ship arrivals is generated in loop 2–8. The loop execution is ended if the next arrival is after the end of time horizon  $H$  or sufficient number of arrivals is created. In Line 7 the next arrival time is calculated.

## 6. Distinctive port features

As shown in the previous sections, ports differ in many ways. The specific features that appear in our data can help to choose the right model for some other port.

One expected port dissimilarity is in the mixture of ship sizes. Different length patterns are expected due to the characteristic location, like the existence of a chain of neighboring ports, or the location in the river estuary. Examples of ship size histograms covering all size classes with the same box ranges are shown in Fig. 10. For example, Los Angeles (Fig. 10(c)) has no small container ships because there are no river waterways nearby and

**Algorithm 3.** DrawReturning( $j, c, L_j, a_{max}$ )

---

```

1: generate first arrival date as proposed in
   Section 4.5.1;  $count \leftarrow 0$ ;
2: while  $date < H$  and  $count < a_{max}$  do
3:   generate number  $x$  according to the continuous
       distribution of ratios  $p_c/L_c$  for the cluster  $c$ 
       (Section 4.4) set processing time  $p = x \times L_j$  for
       call  $(j + count)$ ;
4:   record  $(j + count, L_j, p, date)$ ;
5:    $count \leftarrow count + 1$ ;
6:   choose return time  $\rho$  according to mixture of
       distributions (Section 4.5.1);
7:    $date \leftarrow date + \rho$ ;
8: end while
9: return  $count$ ;

```

---

Table 7. Correlation between  $p_j/L_j$  and  $L_j$ .

| port   | Gdańsk  | Long Beach | Los Angeles | Le Havre  |
|--------|---------|------------|-------------|-----------|
| $r$    | −0.007  | 0.065      | 0.389       | −0.263    |
| $SE_r$ | 0.046   | 0.032      | 0.025       | 0.020     |
| port   | Hamburg | Rotterdam  | Shanghai    | Singapore |
| $r$    | −0.273  | −0.185     | 0.005       | −0.105    |
| $SE_r$ | 0.017   | 0.016      | 0.009       | 0.007     |

even local traffic is oceanic. A similar pattern of ship sizes can be found in Long Beach (not shown here). Conversely, Gdańsk and Shanghai (Fig. 10(b), (d)) are in a chain of local ports, and what is more, the Shanghai port is in large rivers estuaries. Hence, there are large fractions of small ships in the traffic. Though Le Havre is in the Seine estuary it is also the westernmost ocean port in a chain of European ports and hence there is a significant fraction of medium size ([150,300] m) ships for local connections (Fig. 10a). Similar pattern with the domination of medium size ships exist in Singapore (not shown here).

The pattern of ship sizes is connected with the return intervals because small ships operate locally, whereas the largest ones operate in the long ocean lines. This can be verified in Figs. 5(b) and 10(a), (c) for Le Havre, Los Angeles, as well as in Fig. 11 for Rotterdam. In Le Havre there is a broad set of ship size classes and return periods because it is an ocean port, in a chain of local ports. There are almost no short ships and short return times in Los Angeles because it is an ocean port. In Rotterdam (see Fig. 11) ships shorter than 180 m and returning in less than 3 weeks dominate, and very large vessels are not present. This is understandable for an inland port on the banks of Nieuwe Maas (the Maasvlakte terminals traffic is not included into the analyzed data).

Let us conclude the above discussion with a recommendation on suitable STMs for optimization

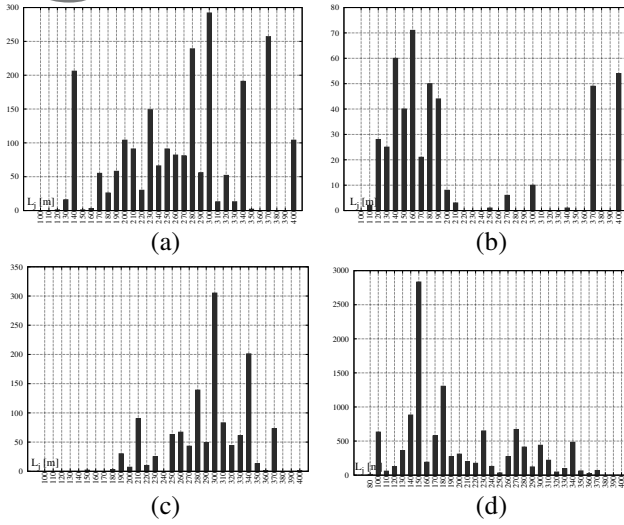


Fig. 10. Examples of ship size histograms with 10 m resolution: Le Havre (a), Gdańsk (b), Los Angeles (c), Shanghai (d).

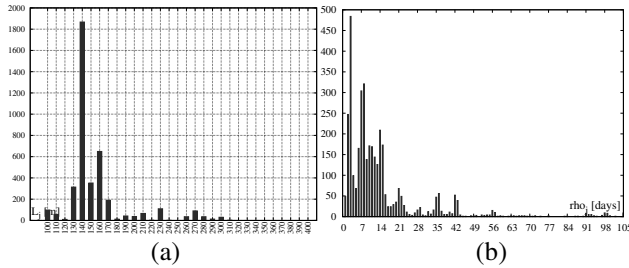


Fig. 11. Rotterdam: histograms of  $L_j$  (a) and return period  $\rho_j$  (b).

problems of port logistics. Location is the key selection determinant. An STM of the most similar port should be used. If the location considered for optimization is an isolated ocean port, then Los Angeles or Long Beach can be chosen. If the port is in a chain of local ports, choose Le Havre or Singapore STMs. If, furthermore, the location is in big river system estuary, then Shanghai STM can be recommended. For an inland port, our Rotterdam STM can be applied.

Another interesting feature is the relationship between the ship length  $L_j$  and its processing time  $p_j$ . It can be expected that, in general, longer ships have longer processing times. In order to compensate for this rather obvious relationship, we decided to investigate the correlation between  $L_j$  and  $p_j/L_j$  (see also Section 4.4). It seems intuitive that such correlation either should not exist (to treat different size classes fairly) or may be negative (because terminal operators would rather dedicate more resources and effort to the largest ships to take advantage of the economy of scale). The coefficients of correlation  $r$  between  $L_j$  and  $p_j/L_j$  and the standard errors  $SE_r$  of these coefficients are shown in Table 7. The correlations in Table 7 are weak according to the

frequently used rule of thumb requiring that  $|r| \geq 0.7$  for a strong correlation. If we use the rule that range  $[r - SE_r, r + SE_r]$  does not include zero, then it can be observed that there are ports where the correlation is nonexistent (Gdańsk, Shanghai) which can be interpreted as a fair treatment of the different classes. There are ports where the correlation is negative as expected. Surprisingly, there are also ports (Los Angeles) where longer ships do take longer time (disproportionately to  $L_j$ ) to be served. A closer look at the terminal management policy is needed to explain these effects.

Let us finish this section by observing that with a growing number of arrivals  $N$  the coefficient of variation of the arrivals calculated over to days of a week and hours of a day decreases (Fig. 9(b); Wawrzyniak et al., 2021) and the feasibility of fitting a continuous distribution as a model of  $L_j$ s for size clusters improves (Figs. 2(c), (d)).

## 7. Conclusions

In this work, we studied container ship traffic patterns in eight ports of the world to develop traffic models for optimization and simulation of port logistics. We developed more advanced and detailed models than those existing previously in the literature. The models respect the relationships between ship sizes, processing times and arrival times. Each size class has its own processing time and arrival time statistical models. Particular attention is paid to return times and their dispersion, but aperiodic ships are also considered. The achieved results provide many general observations, both on the assumptions common to port logistic models, and on generating simulation scenarios (benchmark instances).

Representing real-world phenomena in mathematical models is full of trade-offs. One of our goals was to generalize the observed traffic patterns, but generalizing leads to loss of information. For example, establishing some general model framework is possible, but particular ports require their own variants to adjust not only the distribution numerical parameters but also the actual distribution types (see Table 6). Hence, a trade-off exists between accuracy when keeping details of a given port, and simplicity when using one model with a unique distribution.

Since we decided to build models representing each port separately, a decision-maker should choose between these ports and treat them as examples (archetypes) of ocean ports, ports in local port networks, or ports with local river traffic.

Simulations for alternative terminal designs or creation of test instances for port logistic algorithms are among the many possible uses of this work. Our future work will focus on developing terminal design algorithms and verifying their performances with the help of the models described in this publication.

## Acknowledgment

The authors would like to thank Ronan Kerbiriou from Le Havre University for providing AIS traffic data. This research has been partially supported by a grant from the French National Research Ministry and the PHC Polonium project for the years 2017 and 2018. Éric Sanlaville's work has been partially supported by the Normandy region and the European ERDF through the CLASSE2 project.

## References

- Beasley, J. (2018). *OR-Library*, <http://people.brunel.ac.uk/~mastjjb/jeb/info.html>.
- Bellsola Olba, X., Daamen, W., Vellinga, T. and Hoogendoorn, S.P. (2017). Network capacity estimation of vessel traffic: An approach for port planning, *Journal of Waterway, Port, Coastal, and Ocean Engineering* **143**(5): 04017019.
- Bellsola Olba, X., Daamen, W., Vellinga, T. and Hoogendoorn, S.P. (2018). State-of-the-art of port simulation models for risk and capacity assessment based on the vessel navigational behaviour through the nautical infrastructure, *Journal of Traffic and Transportation Engineering (English Edition)* **5**(5): 335–347.
- Benaglia, T., Chauveau, D., Hunter, D.R. and Young, D.S. (2009). Mixtools: An R package for analyzing mixture models, *Journal of Statistical Software* **32**(6): 1–29.
- Bierwirth, C. and Meisel, F. (2010). A survey of berth allocation and quay crane scheduling problems in container terminals, *European Journal of Operational Research* **202**(3): 615–627.
- Bierwirth, C. and Meisel, F. (2015). A follow-up survey of berth allocation and quay crane scheduling problems in container terminals, *European Journal of Operational Research* **244**(3): 675–689.
- Buhrkal, K., Zuglian, S., Ropke, S., Larsen, J. and Lusby, R. (2011). Models for the discrete berth allocation problem: A computational comparison, *Transportation Research E: Logistics and Transportation Review* **47**(4): 461–473.
- Çağatay, I. and Siu Lee Lam, J. (2021). Optimal energy management and operations planning in seaports with smart grid while harnessing renewable energy under uncertainty, *Omega* **103**: 102445, DOI: 10.1016/j.omega.2021.102445.
- Chen, G. and Yang, Z.-Z. (2014). Methods for estimating vehicle queues at a marine terminal: A computational comparison, *International Journal of Applied Mathematics and Computer Science* **24**(3): 611–619, DOI: 10.2478/amcs-2014-0044.
- Delignette-Muller, M.L. and Dutang, C. (2015). fitdistrplus: An R package for fitting distributions, *Journal of Statistical Software* **64**(4): 1–34.
- Dragovic, B., Park, N.K. and Radmilovic, Z. (2006). Ship-berth link performance evaluation: Simulation and analytical approaches, *Maritime Policy & Management* **33**(3): 281–299.
- Feitelson, D.G., Tsafir, D. and Krakov, D. (2014). Experience with using the parallel workloads archive, *Journal of Parallel and Distributed Computing* **74**(10): 2967–2982.
- Giallombardo, G., Moccia, L., Salani, M. and Vacca, I. (2010). Modeling and solving the tactical berth allocation problem, *Transportation Research B: Methodological* **44**(2): 232–245.
- Gosasang, V., Chandraprakaikul, W. and Kiattisin, S. (2011). A comparison of traditional and neural networks forecasting techniques for container throughput at Bangkok port, *Asian Journal of Shipping and Logistics* **27**(3): 463–482.
- Hedjar, R. and Bounkhe, M. (2019). An automatic collision avoidance algorithm for multiple marine surface vehicles, *International Journal of Applied Mathematics and Computer Science* **29**(4): 759–768, DOI: 10.2478/amcs-2019-0056.
- Imai, A., Yamakawa, Y. and Huang, K. (2014). The strategic berth template problem, *Transportation Research E: Logistics and Transportation Review* **72**: 77–100, DOI: 10.1016/j.tre.2014.09.013.
- Kang, L., Meng, Q. and Tan, K.C. (2020). Tugboat scheduling under ship arrival and tugging process time uncertainty, *Transportation Research E: Logistics and Transportation Review* **144**: 102125, DOI: 10.1016/j.tre.2020.102125.
- Lasdon, L., Fox, R. and Ratner, M. (1974). Nonlinear optimization using the generalized reduced gradient method, *RAIRO—Operations Research—Recherche Opérationnelle* **8**(3): 73–103.
- Li, C., Qi, X. and Song, D. (2016). Real-time schedule recovery in liner shipping service with regular uncertainties and disruption events, *Transportation Research B: Methodological* **93**: 762–788, DOI: 10.1016/j.trb.2015.10.004.
- Liu, C. (2020). Iterative heuristic for simultaneous allocations of berths, quay cranes, and yards under practical situations, *Transportation Research E: Logistics and Transportation Review* **133**: 101814, DOI: 10.1016/j.tre.2019.11.008.
- NEO Research Group (2013). Vehicle routing problem, <https://neo.lcc.uma.es/vrp/>.
- Pachakis, D. and Kiremidjian, A.S. (2003). Ship traffic modeling methodology for ports, *Journal of Waterway, Port, Coastal, and Ocean Engineering* **129**(5): 193–202.
- Reinelt, G. (1995). *TSPLIB*, <http://comopt.ifl.uni-heidelberg.de/software/TSPLIB95/index.html>.
- Schepler, X., Balev, S., Michel, S. and Sanlaville, É. (2017). Global planning in a multi-terminal and multi-modal maritime container port, *Transportation Research E: Logistics and Transportation Review* **100**: 38–62, DOI: 10.1016/j.tre.2016.12.002.
- Shabayek, A. and Yeung, W. (2002). A simulation model for the Kwai Chung container terminals in Hong Kong, *European Journal of Operational Research* **140**(1): 1–11.
- Stahlbock, R. and Voß, S. (2008). Operations research at container terminals: A literature update, *OR Spectrum* **30**(1): 1–52.

- Taillard, E. (1993). Benchmarks for basic scheduling problems, *European Journal of Operational Research* **64**(2): 278–285.
- van Asperen, E., Dekker, R., Polman, M. and de Swaan Arons, H. (2003). Modeling ship arrivals in ports, in S. Chick et al. (Eds), *Proceedings of the 2003 Winter Simulation Conference*, IEEE, New York, pp. 1737–1744.
- Wang, W., Chen, X., Musial, J. and Blazewicz, J. (2020). Two meta-heuristic algorithms for scheduling on unrelated machines with the late work criterion, *International Journal of Applied Mathematics and Computer Science* **30**(3): 573–584, DOI: 10.34768/amcs-2020-0042.
- Wawrzyniak, J., Drozdowski, M. and Sanlaville, É. (2020). Selecting algorithms for large berth allocation problems, *European Journal of Operational Research* **283**(3): 844–862.
- Wawrzyniak, J., Drozdowski, M. and Sanlaville, É. (2021). Container ship traffic model for simulation studies—Additional resources, <http://www.cs.put.poznan.pl/mdrozdowski/stm/>.

**Jakub Wawrzyniak** is the CTO at the TIDK company. He is also a PhD candidate at the Institute of Computer Science of the Poznań University of Technology. His research interests include algorithm design, computational complexity analysis, machine learning, combinatorial optimization, and task scheduling.

**Maciej Drozdowski** is a full professor in the Institute of Computing Science, Poznań University of Technology. His research interests include scheduling for parallel processing, computer performance evaluation, web engineering, design and analysis of algorithms, and combinatorial optimization.

**Éric Sanlaville** is a full professor in the Computer Science Department, Le Havre Normandy University. He is a specialist in operations research applied to scheduling and logistics. His research focuses on decision under uncertainty and dynamic graphs.

Received: 17 November 2021

Revised: 13 April 2022

Re-revised: 14 June 2022

Accepted: 18 July 2022