

Springer Proceedings in Mathematics & Statistics

João Paulo Almeida ·  
Filipe Pereira e Alvelos ·  
Jorge Orestes Cerdeira · Samuel Moniz ·  
Cristina Requejo *Editors*

# Operational Research

IO 2022—OR in Turbulent Times:  
Adaptation and Resilience. XXII Congress  
of APDIO, University of Évora, Portugal,  
November 6–8, 2022

 Springer

# **Springer Proceedings in Mathematics & Statistics**

Volume 437

This book series features volumes composed of selected contributions from workshops and conferences in all areas of current research in mathematics and statistics, including data science, operations research and optimization. In addition to an overall evaluation of the interest, scientific quality, and timeliness of each proposal at the hands of the publisher, individual contributions are all refereed to the high quality standards of leading journals in the field. Thus, this series provides the research community with well-edited, authoritative reports on developments in the most exciting areas of mathematical and statistical research today.

João Paulo Almeida · Filipe Pereira e Alvelos ·  
Jorge Orestes Cerdeira · Samuel Moniz ·  
Cristina Requejo  
Editors

# Operational Research

IO 2022—OR in Turbulent Times: Adaptation  
and Resilience. XXII Congress of APDIO,  
University of Évora, Portugal, November 6–8,  
2022

### *Editors*

João Paulo Almeida  
Department of Mathematics  
Polytechnic Institute of Bragança  
and CeDRI  
Bragança, Portugal

Jorge Orestes Cerdeira  
Department of Mathematics  
Center for Mathematics and Applications  
NOVA School of Science and Technology  
Costa da Caparica, Portugal

Cristina Requejo  
ISEG-Lisbon School of Economics  
and Management  
University of Lisbon  
Lisbon, Portugal

Filipe Pereira e Alvelos  
Algoritmi Research Center and LASI  
University of Minho  
Braga, Portugal

Samuel Moniz  
Department of Mechanical Engineering  
University of Coimbra  
Coimbra, Portugal

ISSN 2194-1009                      ISSN 2194-1017 (electronic)  
Springer Proceedings in Mathematics & Statistics  
ISBN 978-3-031-46438-6              ISBN 978-3-031-46439-3 (eBook)  
<https://doi.org/10.1007/978-3-031-46439-3>

Mathematics Subject Classification: 90-10, 90Bxx, 60-XX, 68U35, 90Cxx, 90-XX

© The Editor(s) (if applicable) and The Author(s), under exclusive license to Springer Nature Switzerland AG 2023

This work is subject to copyright. All rights are solely and exclusively licensed by the Publisher, whether the whole or part of the material is concerned, specifically the rights of translation, reprinting, reuse of illustrations, recitation, broadcasting, reproduction on microfilms or in any other physical way, and transmission or information storage and retrieval, electronic adaptation, computer software, or by similar or dissimilar methodology now known or hereafter developed.

The use of general descriptive names, registered names, trademarks, service marks, etc. in this publication does not imply, even in the absence of a specific statement, that such names are exempt from the relevant protective laws and regulations and therefore free for general use.

The publisher, the authors, and the editors are safe to assume that the advice and information in this book are believed to be true and accurate at the date of publication. Neither the publisher nor the authors or the editors give a warranty, expressed or implied, with respect to the material contained herein or for any errors or omissions that may have been made. The publisher remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

This Springer imprint is published by the registered company Springer Nature Switzerland AG  
The registered company address is: Gewerbestrasse 11, 6330 Cham, Switzerland

Paper in this product is recyclable.

# Foreword

The pages of this book unfold a compendium born from the IO2022—XXII Congress of the Portuguese Society of Operational Research, a pivotal event that convened in the venerable halls of the University of Évora in the year 2022. Guided by the resonant theme “OR in Turbulent Times: Adaptation and Resilience”, the congress provided a dynamic platform for the intersection of ideas, insights and advances in the field of Operational Research (OR). Operational Research, often referred to as the science of decision-making, is a multidisciplinary field that uses quantitative analysis, mathematical modelling and optimisation techniques to unravel complex problems and facilitate informed decisions in a variety of domains. The Congress, a manifestation of scientific dialogue and exchange, captured the essence of OR’s role in navigating the complexities of an ever-evolving world. Exploring the nexus of adaptation and resilience, it delved into the strategies and methodologies that enable industries and societies to respond dynamically to challenges, harnessing the power of data-driven insights to create solutions that withstand the test of change. In essence, the IO2022 Congress embodied the essence of operational research as an intellectual beacon that guides us through the maze of uncertainty, fosters innovation and equips us to thrive even in turbulent times.

Specifically, the collection of 16 scientific papers presented in this book covers a wide range of fields, all underpinned by optimisation and decision support methodologies. Together, these papers contribute to the advancement of various industries and societal domains. From healthcare logistics to supply chain resilience, these papers demonstrate the power of optimisation techniques to address complex real-world challenges. The common thread in these papers is the use of advanced mathematical models, ranging from robust optimisation in radiotherapy planning to multi-objective simulation optimisation in supply chain design. These methods are used to increase efficiency, reduce environmental impact and improve decision-making across sectors. By addressing issues such as transportation, disaster response, manufacturing and energy systems, this compilation underscores the lasting impact of optimisation-driven research on society, fostering smarter, more sustainable practices and systems.

This is the end of Artificial Intelligence's contribution to this preface. Two requests have been made to ChatGPT-3.5:

This book presents a selection of papers presented in the IO2022—XXII Congress of the Portuguese Society of Operational Research, which was run in the University of Évora in 2022, under the motto "OR in Turbulent Times: Adaptation and Resilience". Can you write a paragraph about the congress and what is Operational Research?

I need to write a one-paragraph summary of a book containing 16 scientific papers. The titles and abstracts are given below. I don't want an abstract for each paper. I need an overall summary that highlights the commonalities and contributions to society. Can you do this for me?

and the answers were perfected by DeepL Pro in terms of English usage.

Should we be afraid? I don't think so. These tools are not capable of sensing the challenges that the future will bring to our world and that we, as scientists, must anticipate and address. These tools are not capable of generating truly new, creative and disruptive ideas. Yes, they will be able to contribute to many steps of the research process by analysing literature, summarising data, automatically transforming mathematical expressions into computer code, designing and running computational experiments and so on. They will become ideal research assistants. They will be able to carry out incremental research by combining and rearranging existing methods and models. But this will allow us, human researchers, to focus on the truly creative phases of the research process, on using intuition to decide on research paths, on setting priorities and degrees of importance. Because all artificial intelligence tools depend on computing power, this will always be finite. Individually, each of us is also finite, but together we can go far beyond the limits of any machine. This is what the series of papers collected in this book show us. This is what the sense of community experienced at the IO2022 Congress, and which characterises every initiative of the Portuguese OR Society, shows us. The future is bright!

Porto, Portugal  
August 2023

José Fernando Oliveira

# Acknowledgements

This volume comprises 16 chapters of meticulously selected works presented during IO2022, the XXII Congress of APDIO, the Portuguese Association of Operational Research, hosted at the University of Évora from the 6th to the 8th of November in the year 2022.

Foremost, our acknowledgement extends to the authors of these chapters, for sharing their current research insights, along with their vision in the domain of cutting-edge methodologies and techniques within the sphere of Operational Research, under the topic—Times of Turbulence: Adaptation and Resilience.

We extend a special mention to José Fernando de Oliveira for his contribution as the author of the preface for this compendium.

Sincere appreciation is also directed towards Martin Peters, the Executive Editor for Mathematics, Computational Science, and Engineering at Springer-Verlag, for his suggestions and expert guidance, which have been instrumental throughout the entire duration of this undertaking.

Our deepest gratitude goes to the IO2022 Program Committee members, as well as to Adelaide Cerveira of the University of Trás-os-Montes e Alto Douro, Alcinda Barreiras of the Polytechnic of Porto, Dalila Fontes of the University of Porto, Filipe Gonçalves Rodrigues of the University of Lisbon, Gonçalo Figueira of the University of Porto, Isabel Cristina Correia of NOVA University of Lisbon, Joana Dias of the University of Coimbra, Leonor Santiago Pinto of the University of Lisbon, Lino Costa of the University of Minho, Luís Guimarães of the University of Porto, Manuel Vieira of NOVA University of Lisbon, Marta Mesquita of the University of Lisbon, Miguel Alves Pereira of the University of Lisbon, Nelson Chibeles-Martins of NOVA University of Lisbon, Pedro Coimbra Martins of the Polytechnic Institute of Coimbra and Raquel Bernardino of the University of Lisbon. These esteemed colleagues, who were responsible for meticulously reviewing the presented works, have elevated the scientific level of the chapters featured within this volume. Their excellent work as reviewers is deeply appreciated.

In addition, it is with profound recognition that we acknowledge the indispensable efforts of the Organizing Committee of IO2022, who demonstrated remarkable dedication in planning this memorable congress.

In a closing remark directed to prospective readers, this volume encapsulates 16 chapters that collectively present an overarching panorama of the research performed by the Portuguese Operational Research community. Emphasis is placed on the resolution of practical, intricate challenges and the latest methodological developments in the realm of Operational Research. With these integral components, it is our fervent aspiration that this volume shall stand as a commendable contribution, serving the needs of students, researchers and all individuals engaged in the pursuit of Operational Research excellence.

João Paulo Almeida  
Filipe Pereira e Alvelos  
Jorge Orestes Cerdeira  
Samuel Moniz  
Cristina Requejo

# Contents

|                                                                                                                                     |    |
|-------------------------------------------------------------------------------------------------------------------------------------|----|
| <b>A Study of a State-of-the-Art Algorithm for Mixed-Integer Linear Bilevel Optimization</b> .....                                  | 1  |
| Maria João Alves, Carlos Henggeler Antunes, and Inês Soares                                                                         |    |
| <b>Using Supplier Networks to Handle Returns in Online Marketplaces</b> .....                                                       | 17 |
| Catarina Pinto, Gonçalo Figueira, and Pedro Amorim                                                                                  |    |
| <b>The Multiple Ambulance Type Dispatching and Relocation Problem: Optimization Approaches</b> .....                                | 31 |
| A. S. Carvalho and M. E. Captivo                                                                                                    |    |
| <b>Automated Radiotherapy Treatment Planning Optimization: A Comparison Between Robust Optimization and Adaptive Planning</b> ..... | 49 |
| Hugo Ponte, Humberto Rocha, and Joana Dias                                                                                          |    |
| <b>An Optimization Model for Power Transformer Maintenance</b> .....                                                                | 63 |
| João Dionísio and João Pedro Pedroso                                                                                                |    |
| <b>Multi-objective Finite-Domain Constraint-Based Forest Management</b> .....                                                       | 75 |
| Eduardo Eloy, Vladimir Bushenkov, and Salvador Abreu                                                                                |    |
| <b>Implementation of Geographic Diversity in Resilient Telecommunication Networks</b> .....                                         | 89 |
| Maria Teresa Godinho and Marta Pascoal                                                                                              |    |
| <b>Are the Portuguese Fire Departments Well Located for Fighting Urban Fires?</b> .....                                             | 99 |
| Maria Isabel Gomes, Ana C. Jória, João Pinela, and Nelson Chibeles-Martins                                                          |    |

|                                                                                                                                            |     |
|--------------------------------------------------------------------------------------------------------------------------------------------|-----|
| <b>Parcel Delivery Services: A Sectorization Approach with Simulation</b> .....                                                            | 113 |
| Cristina Lopes, Ana Maria Rodrigues, Elif Ozturk,<br>José Soeiro Ferreira, Ana Catarina Nunes, Pedro Rocha,<br>and Cristina Teles Oliveira |     |
| <b>Decision Support System for Scheduling the Production of Screw Caps in a Flexible Job-Shop Environment</b> .....                        | 125 |
| José Filipe Ferreira and Rui Borges Lopes                                                                                                  |     |
| <b>Bimaterial Three-Dimensional Printing Using Digital Light Processing Projectors</b> .....                                               | 139 |
| Daniel Bandeira and Marta Pascoal                                                                                                          |     |
| <b>Supply Chain Resilience: Tactical-Operational Models, a Literature Review</b> .....                                                     | 157 |
| Márcia Batista, João Pires Ribeiro, and Ana Barbosa-Póvoa                                                                                  |     |
| <b>Applying Deep Learning Techniques to Forecast Purchases in the Portuguese National Health Service</b> .....                             | 179 |
| José Sequeiros, Filipe R. Ramos, Maria Teresa Pereira,<br>Marisa Oliveira, and Lihki Rubio                                                 |     |
| <b>A Three-Level Decision Support Approach Based on Multi-objective Simulation-Optimization and DEA: A Supply Chain Application</b> .....  | 193 |
| Luís Pedro Gomes, António Vieira, Rui Fragoso, Dora Almeida,<br>Luís Coelho, and José Maia Neves                                           |     |
| <b>Modelling Forest Fire Spread Through Discrete Event Simulation</b> .....                                                                | 209 |
| Catarina Santos, Ana Raquel Xambre, Andreia Hall, Helena Alvelos,<br>Susete Marques, Isabel Martins, and Filipe Alvelos                    |     |
| <b>Multiobjective Evolutionary Clustering to Enhance Fault Detection in a PV System</b> .....                                              | 227 |
| Luciana Yamada, Priscila Rampazzo, Felipe Yamada, Luís Guimarães,<br>Armando Leitão, and Flávia Barbosa                                    |     |

# A Study of a State-of-the-Art Algorithm for Mixed-Integer Linear Bilevel Optimization



Maria João Alves, Carlos Henggeler Antunes, and Inês Soares

**Abstract** In this paper, we address mixed-integer linear bilevel optimization problems. In bilevel optimization, a (lower-level) optimization problem is included in the constraints of another (upper-level) optimization problem. Thus, this framework is especially adequate to model hierarchical decision processes. We analyze a state-of-the-art algorithm developed for this type of problems, which is based on an optimal-value-function reformulation and consists of an iterative deterministic bounding procedure. Computational experiments are made with data instances with different characteristics. The performance of the algorithm in the different groups of problems is discussed.

**Keywords** Bilevel optimization · Mixed-integer linear programming · Optimal value-function reformulation

## 1 Introduction

Bilevel optimization (BLO) models are relevant in problems in which decisions are made sequentially, in a non-cooperative manner, by two decision makers (the leader and the follower) controlling different sets of variables to optimize their own objective functions. The leader first sets his/her decision variables (upper-level problem); the follower then optimizes his/her objective function within the options restricted by the leader's decision (lower-level problem), i.e., the follower's problem is a constraint of

---

M. J. Alves (✉)

Faculty of Economics/ INESC Coimbra, University of Coimbra, CeBER, Coimbra, Portugal  
e-mail: [mjalves@fe.uc.pt](mailto:mjalves@fe.uc.pt)

C. H. Antunes

INESC Coimbra, Department of Electrical and Computer Engineering, University of Coimbra, Coimbra, Portugal  
e-mail: [ch@deec.uc.pt](mailto:ch@deec.uc.pt)

I. Soares

INESC Coimbra, Coimbra, Portugal  
e-mail: [inesgsoares@deec.uc.pt](mailto:inesgsoares@deec.uc.pt)

the leader's problem. Since the follower's decision influences the leader's objective function value (and possibly his/her feasible options), the leader must anticipate the (rational) follower's reaction. This type of sequential decision-making process often appears in the management of decentralized organizations and policy making, including in pricing and tariff design problems, energy markets, and defense of critical infrastructures. For instance, in electricity retail markets involving demand response, BLO models allow for capturing the sequence of decisions regarding the price announcement by the electricity retailer (the leader, aiming to design time-of-use tariffs to maximize profit) and the consumers' reaction involving changes in the energy use to minimize cost (subject to quality-of-service constraints). In this paper, we address mixed-integer linear bilevel optimization problems (MILBLO).

Concerning exact procedures for MILBLO problems, the first generic branch-and-bound method was developed by Moore and Bard [1], which was shown to converge when all leader's variables ( $x$ ) are integer or all the variables controlled by the follower ( $y$ ) are continuous. DeNegre and Ralphs [2] built up on the ideas of this method and developed a branch-and-cut algorithm for the case in which all  $x$  and  $y$  variables are integer and the upper-level constraints do not include  $y$  variables. Xu and Wang [3] presented a branch-and-bound algorithm for MILBLO problems where all  $x$  are integer, as well as the values of the functions of  $x$  in the lower-level constraints. The algorithm branches on these functions, generating multiple branches at each node. A branch-and-cut algorithm has been proposed by Caramia and Mari [4] for linear BLO problems with integer  $x$  and  $y$  variables. Fischetti et al. [5] also proposed a branch-and-cut algorithm for MILBLO problems, which extends the valid intersection cuts for MILBLO proposed in [6]. In [5], new BL-specific preprocessing procedures and a general branch-and-cut exact method are developed, whose finite convergence relies on the assumption that continuous  $x$  variables do not appear in the lower-level problem. A different approach has been proposed by Zeng and An [7] for MILBLO problems, allowing continuous and integer variables in both levels. The computation scheme is based on single-level reformulations and decomposition strategies (there is a master problem and subproblems). The reformulation involves the use of KKT conditions, being the master problem a MILP problem with complementarity constraints. Bolusani and Ralphs [8] propose a framework for MILBLO based on a generalization of Benders' decomposition. The authors also provide a timeline of the main developments in the evolution of exact algorithms for MILBLO up to 2020, indicating the types of variables (continuous, binary, and general integer) that are supported in the upper and lower levels.

The works of Mitsos [9] and Lozano and Smith [10] are representative of approaches that use optimal-value-function reformulations of the lower-level problem (see Sect. 2). Mitsos [9] proposed a deterministic algorithm for general mixed-integer nonlinear BLO problems to obtain a sequence of increasingly tighter upper and lower bounds for the upper-level objective function. The algorithm finishes when the difference between these bounds is below a given tolerance  $\epsilon$ , ensuring

$\epsilon$ -convergence of the algorithm. Lozano and Smith [10] address MILBLO problems in which all  $x$  are integer and the functions of these variables in the lower-level constraints are also integer-valued, profiting from this assumption to strengthen the formulation and proving convergence to a global optimal solution. Both algorithms in [9] and [10] solve a relaxation of the BLO problem with disjunctive constraints generated from optimal follower's responses (already known) to obtain upper bounds for the maximizing upper-level objective function. Bilevel feasible solutions are used to obtain lower bounds. The authors do not prescribe tailored techniques for solving the subproblems that arise in their procedures, assuming that these subproblems are solved by appropriate algorithms available in the literature. In the case of the Lozano-Smith algorithm, all the subproblems are MILP problems that can be solved by a MIP solver (e.g. cplex, gurobi, etc.). The computational experiments reported in [10] showed that Lozano-Smith algorithm outperformed the state-of-the-art algorithm by Xu and Wang [3] for the testbed provided in [3]. The Lozano-Smith algorithm has been used to develop specialized algorithms for optimizing dynamic electricity tariffs [11] and for optimizing pricing strategies at electric vehicle charging stations considering competition from other stations [12]. In both cases, the proposed algorithms benefit from the special structure of the problems.

Tavashoğlu et al. [13] proposed a value-function reformulation to deal with BLO problems with multiple followers, in which the upper-level problem is pure integer and the lower-level problems are mixed-integer. This reformulation uses generalized value functions proposed by the authors for the followers' mixed-integer programs, and takes advantage of the integrality of  $x$  and the corresponding constraint matrix in the followers' problems to enumerate solutions.

In a recent comprehensive survey article, Kleinert et al. [14] review tailored approaches exploiting mixed-integer programming techniques to solve BLO problems. In addition to bilevel problems with convex and, in particular, linear lower-level problems, modern algorithmic approaches to solve mixed-integer bilevel problems that contain integrality constraints in the lower level are presented.

In this paper, we present a study of the state-of-the-art Lozano-Smith algorithm, which is based on the optimal-value reformulation concept for MILBLO. The fundamentals of BLO and, in particular, for MILBLO, are presented in Sect. 2. The Lozano-Smith algorithm is described in Sect. 3; in Sect. 3.1, the algorithm is particularized for a special class of MILBLO problems, in which the upper-level variables do not appear in the lower-level constraints. Computational experiments are presented in Sect. 4, considering different features of the algorithm and instances with larger feasible regions. Conclusions are drawn in Sect. 5.

## 2 Fundamentals of Bilevel Optimization

A general BLO problem can be formulated as follows:

$$\begin{aligned}
 & \max_{x \in X, y} F(x, y) \\
 & s.t. \quad G(x, y) \leq 0 \\
 & \quad \max_{y \in Y} f(x, y) \\
 & \quad s.t. \quad g(x, y) \leq 0
 \end{aligned}$$

where  $x \in \mathbb{R}^{n_1}$  is the vector of variables controlled by the *leader*,  $y \in \mathbb{R}^{n_2}$  is the vector of variables controlled by the *follower*, and  $X$  and  $Y$  are compact sets.

Let  $\Psi(x)$  denote the set of optimal solutions of the lower-level problem parameterized on  $x$ , which is called the follower's *rational reaction set* to a given  $x$ :

$$\Psi(x) := \arg \max_{y \in Y} \{f(x, y) : g(x, y) \leq 0\}$$

Let  $S$  be the set of all constraints of the problem:

$$S = \{(x, y) \in X \times Y : G(x, y) \leq 0, g(x, y) \leq 0\}$$

The BLO problem consists of optimizing the upper-level objective function  $F(x, y)$  over the feasible set, which is generally called *induced region*:

$$IR = \{(x, y) : (x, y) \in S, y \in \Psi(x)\}$$

We consider the optimistic formulation of the BLO problem. In this case, the leader also optimizes over the lower-level outcome  $y \in \Psi(x)$  if the lower-level solution set  $\Psi(x)$  is not a singleton. That is, if the lower-level problem has multiple optimal solutions for a given upper-level choice  $x$ , then the upper-level optimizer chooses the best one for the leader.

We can also write the BLO problem as the *optimal-value-function* reformulation:

$$\begin{aligned}
 & \max_{x \in X, y \in Y} F(x, y) \\
 & s.t. \quad G(x, y) \leq 0 \\
 & \quad g(x, y) \leq 0 \\
 & \quad f(x, y) \geq \phi(x)
 \end{aligned}$$

where  $\phi(x)$  is the so-called optimal-value-function of the lower-level problem:

$$\phi(x) := \max_{y \in Y} \{f(x, y) : g(x, y) \leq 0\}.$$

Solving a BLO problem is very challenging due to its intrinsic nonconvexity and nondifferentiability [15]. Even the linear bilevel problem has been shown to be strongly NP-hard [16]. The existence of discrete variables, particularly in the lower-level problem, further heightens the difficulty of solving the BLO problem. Thus, as will be discussed below, it is important to make the most of the structure and features of the BLO model to develop computationally efficient algorithms.

Let us consider BLO problems with linear objective functions and constraints, including continuous and/or discrete variables. The mixed-integer linear bilevel optimization (MILBLO) problem can be formulated as follows:

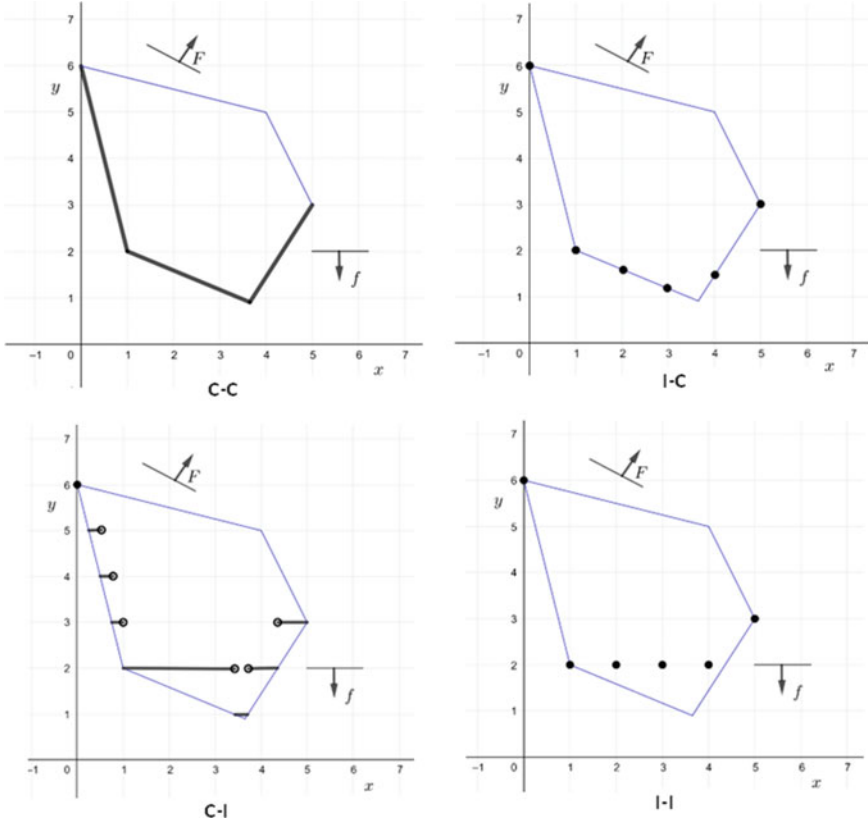
$$\begin{aligned}
 & \max_{x,y} F(x, y) = c_1x + d_1y \\
 & \text{s.t. } A_1x + B_1y \leq b_1 \\
 & \quad x \in X \\
 & \max_y f(x, y) = c_2x + d_2y \\
 & \text{s.t. } A_2x + B_2y \leq b_2 \\
 & \quad y \in Y
 \end{aligned} \tag{1}$$

with  $c_1, c_2 \in \mathbb{R}^{n_1}$ ,  $d_1, d_2 \in \mathbb{R}^{n_2}$ ,  $A_1 \in \mathbb{R}^{m_1 \times n_1}$ ,  $A_2 \in \mathbb{R}^{m_2 \times n_1}$ ,  $B_1 \in \mathbb{R}^{m_1 \times n_2}$ ,  $B_2 \in \mathbb{R}^{m_2 \times n_2}$  and  $b_1 \in \mathbb{R}^{m_1}$ ,  $b_2 \in \mathbb{R}^{m_2}$ . The sets  $X \subseteq \mathbb{R}^{n_1}$  and  $Y \subseteq \mathbb{R}^{n_2}$  may include upper and lower bounds on the variables and impose integrality constraints on all or some of the variables  $x$  and  $y$ , respectively.

In addition to the linear BLO problem with continuous variables in the upper and lower-level problems (C-C case), three cases can be distinguished considering that all variables at one level are continuous or integer: only the upper-level variables are integer (I-C case); both the upper and the lower-level variables are integer (I-I case); only the lower-level variables are integer (C-I case).

In the C-C and I-C cases, a common approach to solve the problem is replacing the lower-level problem with its Karush-Kuhn-Tucker (KKT) conditions, which provide the necessary and sufficient conditions for optimality. The resulting single-level optimization model is nonlinear due to the complementarity conditions, but these conditions can be linearized by means of auxiliary binary variables and constraints resulting in a mixed-integer linear programming (MILP) model, which can then be solved using an off-the-shelf MIP solver.

If the constraint set  $S$  is non-empty and compact, and there are no upper-level constraints involving both the upper-level and the lower-level variables (usually called connecting or coupling constraints), then the linear BLO problem (C-C case) has certainly an optimal solution. However, if there are connecting constraints, the existence of a solution is not guaranteed and more assumptions are needed [17]. Vicente et al. [18] showed that the I-C and I-I problems always have optimal solutions under hypotheses similar to the C-C case. The existence of an optimal solution to the C-I problem seems to be a more difficult task: the problem may not have an optimal solution even if there are no connecting constraints, because the induced region may not be closed.



**Fig. 1** Induced regions of problems C-C, I-C, C-I, I-I

Figure 1 illustrates induced regions of problems C-C, I-C, C-I and I-I. The following property was proved in [18]: the induced regions of I-C and I-I problems are included in the induced regions of the C-C and C-I problems, respectively. This property can be observed in Fig. 1.

### 3 The Lozano-Smith Algorithm for MILBLO Problems

The Lozano-Smith algorithm [10] is devoted to the MILBLO problem, requiring all upper-level variables to be integer. It also assumes that  $A_2x$  in the MILBLO formulation (1) is integer-valued for all  $x \in X$ , which guarantees the exactness of the algorithm.

The single-level optimization problem obtained by relaxing the optimality requirement of the follower's objective function is the so-called *high point relaxation* (HPR) problem. This problem consists of computing the point  $(x, y)$  satisfying the leader's and the follower's constraints that optimizes the leader's objective function:

$$\begin{aligned} \max_{x,y} F(x, y) &= c_1x + d_1y \\ \text{s.t. } (x, y) &\in S \end{aligned} \quad (2)$$

$S$  is the constraint set, which is defined for the MILBLO problem (1) as:  $S = \{(x, y) \in X \times Y : A_1x + B_1y \leq b_1, A_2x + B_2y \leq b_2\}$ .

An optimal solution to the HPR problem gives a valid upper bound for  $F^*$  (the optimal upper-level objective function value for the bilevel problem).

Let  $Y(x) = \{y \in Y : B_2y \leq b_2 - A_2x\}$ .

A solution  $(x, y) \in S$  is feasible to the MILBLO problem if and only if  $f(x, y) \geq f(x, \hat{y})$  for every  $\hat{y} \in Y(x)$ .

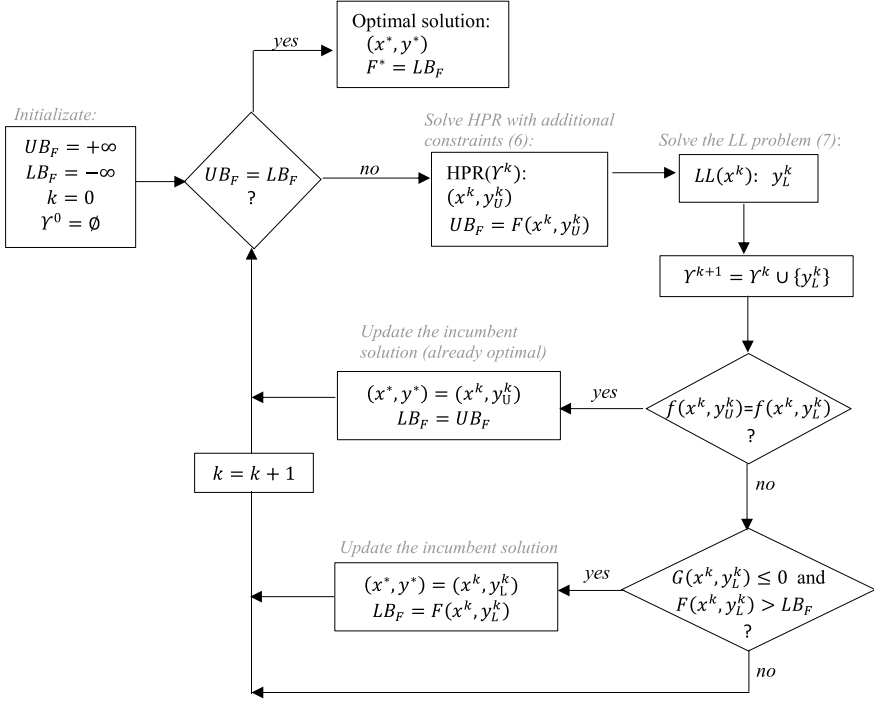
This proposition means that disjunctive constraints can be added to the HPR problem (in which the leader controls  $x$  and  $y$  variables) to enforce bilevel feasibility. Let  $\Upsilon = \bigcup_{x \in \chi} Y(x)$ , with  $\chi = \{x \in X : \exists y \text{ such that } (x, y) \in S\}$ . These disjunctive constraints impose that for every  $\hat{y} \in \Upsilon$ , the leader must either select an  $(x, y)$  such that  $f(x, y) \geq f(x, \hat{y})$ , or *block* the follower's solution  $\hat{y}$  by selecting an  $x$  such that  $\hat{y} \notin Y(x)$  (i.e., at least one of the lower-level constraints is not satisfied).

Therefore, the additional constraints to be added to the HPR problem (2) are:

$$\begin{aligned} &\hat{y} \notin Y(x) \vee f(x, y) \geq f(x, \hat{y}), \quad \forall \hat{y} \in \Upsilon \\ \iff &(A_{2j}x + B_{2j}\hat{y} > b_{2j} \text{ for at least one } j \in \{1 \dots m_2\}) \vee \\ &(c_2x + d_2y \geq c_2x + d_2\hat{y}), \quad \forall \hat{y} \in \Upsilon \end{aligned} \quad (3)$$

The HPR (2) with the additional constraints (3) is equivalent to the bilevel problem (1). Thus, if all the follower's feasible solutions in  $\Upsilon$  could be computed, solving (2) with (3) would give the optimistic optimal solution  $(x, y)$  for the MILBLO problem. However, computing all the follower's feasible solutions is not practicable. Therefore, a sub-set  $\Upsilon^k = \{y^1, \dots, y^k\} \subset \Upsilon$  is considered, leading to an upper bound for the optimal  $F$  value ( $F^*$ ). A sequence of HPR problems (2) including constraints (3) is iteratively solved by adding to  $\Upsilon^k$  a follower's feasible solution  $y^{k+1}$  in each iteration to obtain  $\Upsilon^{k+1}$  (thus building a sample of  $\Upsilon$ ), providing successively tighter upper bounds for  $F^*$ .

In the first iteration of the algorithm, the HPR problem (2) is solved, i.e., an empty set  $\Upsilon^0$  is considered. The lower-level problem is then solved for the resulting solution  $x$ , giving a solution  $y$  that is included into  $\Upsilon^0$  to form  $\Upsilon^1$ ; the HPR problem with additional constraints is solved again considering the updated set  $\Upsilon^1$ . The steps of the Lozano-Smith algorithm are summarized in the flowchart in Fig. 2 and described below.



**Fig. 2** General scheme of the Lozano-Smith algorithm

In this algorithm, solving  $\text{HPR}(\Upsilon^k)$  means solving a MILP problem resulting from (2) with the logical conditions (3) for  $\Upsilon^k$  converted into linear constraints with additional binary variables. Let us look at the first term in (3) for a given solution  $y^\nu \in \Upsilon^k$  (which contains  $k$  solutions). Since there are  $m_2$  constraints in the lower-level problem, then these conditions become  $A_{2j}x + B_{2j}y^\nu > b_{2j}$  for at least one  $j \in 1 \dots m_2$ . By assumption,  $A_{2j}x$  is integer-valued for any  $x$ , so conditions (3) for  $\Upsilon^k$  can be equivalently written as:

$$(A_{2j}x \geq \eta_j^\nu \text{ for at least one } j \in \{1 \dots m_2\}) \vee f(x, y) \geq f(x, y^\nu), \forall y^\nu \in \Upsilon^k \quad (4)$$

with  $\eta_j^\nu = \lfloor b_{2j} - B_{2j}y^\nu \rfloor + 1$ ,  $j = 1 \dots m_2$ .

Condition (4) can be reformulated as MILP constraints with auxiliary variables. Let  $w_j^\nu$ ,  $j = 1 \dots m_2$ ,  $\nu = 1 \dots k$  be binary variables such that  $w_j^\nu = 1$  if constraint  $j$  blocks solution  $y^\nu$ , i.e., constraint  $A_{2j}x \geq \eta_j^\nu$  holds; if  $w_j^\nu = 0$ , then constraint  $j$  may or may not block  $y^\nu$ . Thus,  $\text{HPR}(\Upsilon^k)$  can be formulated as:

$$\begin{aligned} \max_{x,y} \quad & F(x, y) = c_1 x + d_1 y \\ \text{s.t.} \quad & \end{aligned} \quad (5)$$

$$A_{2j}x \geq -M_j^1 + (M_j^1 + \eta_j^v)w_j^v, \quad j = 1 \dots m_2, \quad \forall y^v \in \Upsilon^k \quad (5a)$$

$$f(x, y) \geq f(x, y^v) - M_v^2 \sum_{j=1}^{m_2} w_j^v, \quad \forall y^v \in \Upsilon^k \quad (5b)$$

$$(x, y) \in S$$

$$w_j^v \in \{0, 1\}, \quad j = 1 \dots m_2, \quad v = 1 \dots k$$

The first group of constraints (5a) ensures that, for every lower-level constraint  $j$ ,  $w_j^v = 0$  for all  $y^v \in \Upsilon^k$  that satisfy constraint  $j$  for that  $x$ , i.e.  $A_{2j}x < \eta_j^v$ . If  $y^v$  satisfies all constraints  $j = 1 \dots m_2$ , it means that  $y^v$  is not blocked and, so, the second condition in (4) must be satisfied. In this case, since  $w_j^v = 0$  for all  $j = 1 \dots m_2$ , then  $M_v^2 \sum_{j=1}^{m_2} w_j^v = 0$  and (5b) ensures that  $f(x, y) \geq f(x, y^v)$ . If  $w_j^v = 1$  for some  $v$ , then  $A_{2j}x \geq \eta_j^v$  which means that  $y^v$  is blocked by constraint  $j$ . In this case, the condition  $f(x, y) \geq f(x, y^v)$  does not need to be satisfied; the constraint for  $y^v$  in (5b) becomes redundant.

The authors [10] consider a more compact formulation of  $\text{HPR}(\Upsilon^k)$ , which reduces the number of constraints of (5a) from  $m_2 \times k$  to  $m_2$ . For each solution  $y^v \in \Upsilon^k$ , the following set is defined:

$$\mathcal{B}(y^v) = \left\{ (y^{v'}, j) \mid \eta_j^{v'} \geq \eta_j^v, \quad y^{v'} \in \Upsilon^k, \quad j = 1 \dots m_2 \right\}$$

This set represents all ordered pairs  $(y^{v'}, j)$  such that if  $x$  blocks solution  $y^{v'}$  by constraint  $j$ , then  $x$  also blocks solution  $y^v$  by constraint  $j$ . Consider, for instance, three solutions  $y^{v1}$ ,  $y^{v2}$  and  $y^{v3}$  such that  $\eta_j^{v1} \geq \eta_j^{v2} \geq \eta_j^{v3}$ . Then,  $\mathcal{B}(y^{v1})$  contains  $\{(y^{v1}, j)\}$ ,  $\mathcal{B}(y^{v2})$  contains  $\{(y^{v1}, j), (y^{v2}, j)\}$  and  $\mathcal{B}(y^{v3})$  contains  $\{(y^{v1}, j), (y^{v2}, j), (y^{v3}, j)\}$ . The main idea of the compact formulation is that  $w_j^{v1} = 1$  if constraint  $j$  blocks solutions  $y^{v1}$ ,  $y^{v2}$ ,  $y^{v3}$ , but  $w_j^{v2} = w_j^{v3} = 0$ .

Thus, the alternative formulation to (5) for the  $\text{HPR}(\Upsilon^k)$ , which is the one implemented in the algorithm, is:

$$\begin{aligned} \max_{x,y} \quad & F(x, y) = c_1 x + d_1 y \\ \text{s.t.} \quad & \end{aligned} \quad (6)$$

$$A_{2j}x \geq -M_j^1 + \sum_{y^v \in \Upsilon^k} (M_j^1 + \eta_j^v)w_j^v, \quad j = 1 \dots m_2 \quad (6a)$$

$$f(x, y) \geq f(x, y^v) - M_v^2 \sum_{(y^{v'}, j) \in \mathcal{B}(y^v)} w_j^{v'}, \quad \forall y^v \in \Upsilon^k \quad (6b)$$

$$(x, y) \in S$$

$$w_j^v \in \{0, 1\}, \quad j = 1 \dots m_2, \quad v = 1 \dots k$$

Since  $w_j^v$  can be equal to 0 if constraint  $j$  blocks or does not block  $y^v$ , then  $w_j^v = 0$  by default and  $w_j^v = 1$  only when necessary. This implies that at most one  $w_j^v$  will be equal to 1 for each  $j = 1 \dots m_2$  and  $\sum_v w_j^v \leq 1$ . Therefore, constraints (6a) are never made impossible by some  $+M_j^1$  within the summation that is not canceled out by the first  $-M_j^1$ .

Let us suppose that  $y^{v1}$  is blocked by constraint  $j$  and  $w_j^{v1} = 1$ . Constraint  $j$  in (6a) becomes  $A_{2j}x \geq \eta_j^{v1}$ , and the constraints in (6b) must be redundant for  $y^{v1}$  and for all  $y^{va}$  such that  $\eta_j^{va} \leq \eta_j^{v1}$  (interpret  $va$  as representing the index of all solutions under these conditions). Since  $w_j^{v1} = 1$  and  $w_j^{va} = 0$  for all  $y^{va}$ , then for each of these solutions  $\sum_{(y^{v'}, j) \in \mathcal{B}(y^v)} w_j^{v'} = 1$  and the corresponding constraint in (6b) becomes  $f(x, y) \geq f(x, y^v) - M_v^2$  (i.e., a redundant constraint).

Now, let us consider the case of  $y^v$  that is not blocked by any constraint  $j$ . So,  $w_j^v = 0$  for all  $j = 1 \dots m_2$ . Since  $A_{2j}x < \eta_j^v, \forall j = 1 \dots m_2$ , then  $A_{2j}x < \eta_j^{v'}$  for all  $y^{v'}$  such that  $\eta_j^{v'} \geq \eta_j^v$ , which means that  $w_j^{v'} = 0$  for all  $(y^{v'}, j) \in \mathcal{B}(y^v)$ . Thus,  $f(x, y) \geq f(x, y^v)$  must be satisfied in (6b).

Indications on how to set large enough  $M_j^1$  and  $M_v^2$  values required in (6a) and (6b) are given in [10].

In iteration  $k$  of the algorithm, let  $(x^k, y_U^k)$  be the optimal solution obtained from solving HPR( $\Upsilon^k$ ) (6). Then,  $UB_F = F(x^k, y_U^k)$  is the current upper bound for the upper-level objective function. The lower-level problem for the leader's solution  $x^k$ , i.e.,  $LL(x^k)$ , is then solved:

$$\begin{aligned} \max_y \quad & f(x^k, y) = c_2 x^k + d_2 y \\ \text{s.t.} \quad & B_2 y \leq b_2 - A_2 x^k \\ & y \in Y \end{aligned} \tag{7}$$

Let the optimal follower's reaction to the leader's solution  $x^k$  be  $y_L^k$ . This solution is inserted into the set of lower-level solutions for the next iteration:  $\Upsilon^{k+1} = \Upsilon^k \cup \{y_L^k\}$ . If, for the same leader's solution  $x^k$ , the lower-level solutions  $y_U^k$  and  $y_L^k$  are such that  $f(x^k, y_U^k) = f(x^k, y_L^k)$ , then no other feasible solution exists that is better than  $(x^k, y_U^k)$  for the leader,  $LB_F = UB_F$ , and the optimal solution to the MILBLO problem is found. Otherwise, the upper-level feasibility should be checked for the solution  $(x^k, y_L^k)$  since lower-level decision variables may appear in the upper-level constraints. If upper-level feasibility is satisfied and the upper-level objective function value for  $(x^k, y_L^k)$  is better than the current  $LB_F$ , then this solution becomes the incumbent one and the lower bound is updated  $LB_F = F(x^k, y_L^k)$ . If  $LB_F = UB_F$  then the algorithm stops. Otherwise, the next HPR( $\Upsilon^{k+1}$ ) is solved.

Lozano and Smith [10] propose a sampling scheme that leverages information obtained in the nodes of the branch-and-bound tree used to solve the HPR problem (2) to generate an initial set of solutions  $\Upsilon^0$ . This scheme analyzes every node with a feasible HPR solution  $(\hat{x}, \hat{y}_U)$ . If  $\hat{y}_U$  optimizes  $LL(\hat{x})$  then  $(\hat{x}, \hat{y}_U)$  is bilevel feasible and  $\hat{y}_U$  is inserted into the sample. Otherwise, the optimal follower's response  $\hat{y}_L$  for

that  $\hat{x}$  is inserted into the sample. In both cases, the  $LB_F$  is updated if a better value for  $F$  is obtained and the solution is bilevel feasible (note that this may not happen with  $(\hat{x}, \hat{y}_L)$  because it may not satisfy the upper-level constraints). The sampling procedure stops when the best upper bound from the branch-and-bound tree is equal to  $LB_F$  or a maximum initial sample size limit is reached. When this sampling scheme is implemented, the first HPR problem in the cycle already considers all the solutions included in the initial sample. The authors propose additional techniques to strengthen the  $\text{HPR}(\Upsilon^k)$  formulation, namely fixing variables and supervalid inequalities. In the computational experiments in Sect. 4, these additional techniques are not considered.

### 3.1 Solving a Special Class of MILBLO Problems

A special class of MILBLO problems is when upper-level variables do not appear in lower-level constraints and lower-level variables do not appear in upper-level constraints. These conditions are common in pricing problems, in which the leader's decision variables  $x$  are prices that are charged to consumers. For example, in electricity retail markets, the retailer sets electricity prices and consumers respond by trying to minimize their electricity bill. Lower-level decision variables and constraints model the operation of appliances, as consumers can reschedule the working periods, change thermostat settings, etc., in order to minimize costs. In these problems, the prices  $x$  appear in the objective function of the lower-level problem but generally do not appear in the constraints of this problem. Also, the consumer variables  $y$  do not enter the leader's constraints, because these usually concern only price conditions. In pricing problems, both objective functions involve products of  $x$  and  $y$  variables (typically prices and quantities). When  $y$  variables are implicitly defined by binary control variables, as in [11], or they are integer, then the products  $x \cdot y$  can be linearized leading to a model with the following structure:

$$\begin{aligned}
 & \max_{x,y} F(x, y) = c_1x + d_1y \\
 & \text{s.t. } A_1x \leq b_1 \\
 & \quad x \in X \\
 & \max_y f(x, y) = c_2x + d_2y \\
 & \text{s.t. } B_2y \leq b_2 \\
 & \quad y \in Y
 \end{aligned} \tag{8}$$

In this setting, any optimal solution  $y_L^k$  to the lower-level problem for a given  $x^k$  is feasible to the BLO problem for any feasible  $x$ . This enables to simplify the Lozano-Smith algorithm (Fig. 2) because: (i) the test of whether the upper-level constraints are satisfied for the solution  $(x^k, y_L^k)$  does not need to be done; (ii) the additional constraints in  $\text{HPR}(\Upsilon^k)$  are much simpler, not requiring the additional binary variables  $w_j^y$ . That is, since  $y_L^k \in Y(x)$  for any  $x$ , the first condition in the

disjunctive constraints in (3) is never satisfied, so the second condition, i.e.  $f(x, y) \geq f(x, y_L^k)$ , must be ensured in  $\text{HPR}(\Upsilon^k)$ . The problem  $\text{HPR}(\Upsilon^k)$  for model (8) takes the following form:

$$\begin{aligned} \max_{x,y} \quad & F(x, y) = c_1x + d_1y \\ \text{s.t.} \quad & f(x, y) \geq f(x, y^v) \quad , \quad \forall y^v \in \Upsilon^k \\ & (x, y) \in S \end{aligned} \tag{9}$$

Note that the simplification (ii) is the most important one to streamline the algorithm, which results from upper-level variables not appearing in lower-level constraints.

## 4 Computational Results

The Lozano-Smith algorithm presented in Sect. 3 for generic MILBLO problems and a special class of MILBLO problems (cf. Sect. 3.1) has been used to solve a set of 10 problems, one for each  $n_1 = n_2$  and  $m_1 = m_2 = 0.4n_1$  of their paper [10]. The code of the problems is  $n_1\_1$ , with  $n_1$  taking values in 10, 60, 110, 160, 210, 260, 310, 360, 410, 460, and  $\_1$  meaning the first instance of group  $n_1$ .

Four sets of experiments have been performed using the Java code also used in [10]:

- Exp. 1—using the original version of the code as proposed by [10];
- Exp. 2—using the original version but disabling the construction of the initial sample;
- Exp. 3—using the original version but considering all the right-hand sides of the constraints multiplied by 10, thus enlarging the feasible region;
- Exp. 4—modifying the instances of Exp. 3 into the special class of MILBLO problems (8), so that lower-level decision variables do not appear in upper-level constraints and vice-versa, and disabling the construction of the initial sample.

The aim is to assess the algorithm performance, the importance of building a good quality initial sample, and make a critical assessment of the difficulty of convergence in practice when the feasible region contains a high number of potential solutions.

A computation budget of 1800 s was given to solve each instance of Exp. 1, 2 and 4, and 3600 s for Exp. 3, including a maximum of 900 s for building the initial sample in Exp. 1 and 3. All the runs were carried out in a computer with an Intel Core i7-7700K CPU@3.6GHz and 64GB RAM. The instance data are available at [https://home.deec.uc.pt/~ch/data\\_MILBLO](https://home.deec.uc.pt/~ch/data_MILBLO).

**Table 1** Results of Exp. 1

| Instance     | Initial Sample |         |          | Algorithm |         |              |                 | Run time (s) |
|--------------|----------------|---------|----------|-----------|---------|--------------|-----------------|--------------|
|              | $UB_F$         | $LB_F$  | # sample | # iter.   | $UB_F$  | $LB_F = F^*$ | $\# \Upsilon^k$ |              |
| <i>10_I</i>  | 399.286        | 351.833 | 1        | 1         | 351.833 | 351.833      | 1               | 0.17         |
| <i>60_I</i>  | 196.289        | 153.197 | 79       | 1         | 153.197 | 153.197      | 1               | 26.17        |
| <i>110_I</i> | 205.882        | 181.667 | 27       | 1         | 181.667 | 181.667      | 27              | 2.33         |
| <i>160_I</i> | 186.625        | 165.000 | 23       | 1         | 165.000 | 165.000      | 23              | 4.29         |
| <i>210_I</i> | 146.766        | 136.800 | 23       | 1         | 136.800 | 136.800      | 23              | 4.46         |
| <i>260_I</i> | 160.463        | 139.000 | 12       | 1         | 139.000 | 139.000      | 12              | 9.09         |
| <i>310_I</i> | 142.763        | 117.000 | 74       | 1         | 117.000 | 117.000      | 74              | 69.47        |
| <i>360_I</i> | 166.088        | 133.000 | 170      | 1         | 133.000 | 133.000      | 170             | 412.58       |
| <i>410_I</i> | 139.004        | 103.504 | 32       | 1         | 103.504 | 103.504      | 32              | 120.85       |
| <i>460_I</i> | 154.344        | 97.588  | 222      | 1         | 97.588  | 97.588       | 222             | 465.00       |

**Table 2** Results of Exp 2

| Instance     | Initial iteration ( $k = 0$ ) |           | Algorithm                |         |         |         | Run time (s) |
|--------------|-------------------------------|-----------|--------------------------|---------|---------|---------|--------------|
|              | $UB_F$                        | $LB_F$    | # iter.= $\# \Upsilon^k$ | $UB_F$  | $LB_F$  | $F^*$   |              |
| <i>10_I</i>  | 399.286                       | 351.833   | 1                        | 351.833 | 351.833 | 351.833 | 0.15         |
| <i>60_I</i>  | 196.289                       | 120.706   | 35                       | 153.197 | 153.197 | 153.197 | 20.39        |
| <i>110_I</i> | 205.882                       | $-\infty$ | 5                        | 181.667 | 181.667 | 181.667 | 1.76         |
| <i>160_I</i> | 186.625                       | 37.953    | 8                        | 165.000 | 165.000 | 165.000 | 5.20         |
| <i>210_I</i> | 146.766                       | 2.003     | 4                        | 136.800 | 136.800 | 136.800 | 2.57         |
| <i>260_I</i> | 160.463                       | 98.245    | 5                        | 139.000 | 139.000 | 139.000 | 5.01         |
| <i>310_I</i> | 142.763                       | 60.485    | 16                       | 117.000 | 117.000 | 117.000 | 22.98        |
| <i>360_I</i> | 166.088                       | 0.328     | 52                       | 133.000 | 133.000 | 133.000 | 468.44       |
| <i>410_I</i> | 139.004                       | -44.908   | 20                       | 103.504 | 103.504 | 103.504 | 69.26        |
| <i>460_I</i> | 154.344                       | -59.439   | 78                       | 100.181 | 97.588  | —       | §            |

§—no solution after 1800s

The results of experiments 1-4 are displayed in Tables 1, 2, 3 and 4, respectively.

Table 1 shows the results of Exp. 1, considering the original version of the algorithm including the construction of the initial sample. The first column identifies the instance (the same identifier that is used in [10]). The next three columns under the header *Initial Sample* refer to the upper-bound and lower-bound of  $F$  in iteration  $k = 0$ , and the size of the initial sample (*# sample*) to be used in iteration 1 of the algorithm: the  $UB_F$  results from the HPR without additional constraints (problem (2)) and the  $LB_F$  results from the bilevel feasible solutions in the initial sample. Using this sample, the algorithm finds the optimal solution in just one iteration. The other columns present the number of iterations (*# iter.*), the upper-bound and the lower-bound of  $F$  at the end of the algorithm ( $UB_F$  and  $LB_F$ , respectively),

**Table 3** Results of Exp 3

| Instance      | Initial Sample |          |          | Algorithm |          |                 |                 | Run time (s) |
|---------------|----------------|----------|----------|-----------|----------|-----------------|-----------------|--------------|
|               | $UB_F$         | $LB_F$   | # sample | # iter.   | $UB_F$   | $LB_F$          | $\# \Upsilon^k$ |              |
| <i>10_1a</i>  | 3995.429       | 3914.333 | 38       | 6         | 3914.333 | 3914.333= $F^*$ | 43              | 1.93         |
| <i>60_1a</i>  | 1999.511       | 1554.536 | 3823     | 1         | —        | —               | 3823            | §            |
| <i>110_1a</i> | 2130.228       | 1994.291 | 6377     | 1         | —        | —               | 6377            | §            |
| <i>160_1a</i> | 1978.199       | 1439.700 | 4527     | 1         | —        | —               | 4527            | §            |
| <i>210_1a</i> | 1706.102       | 1539.000 | 4464     | 1         | —        | —               | 4464            | §            |
| <i>260_1a</i> | 1669.148       | 1540.380 | 3195     | 1         | —        | —               | 3195            | §            |
| <i>310_1a</i> | 1639.361       | 1021.809 | 1740     | 1         | 1620.601 | 1023.727        | 1740            | 2411.43      |
| <i>360_1a</i> | 1775.679       | 1190.934 | 767      | 1         | —        | —               | 767             | §            |
| <i>410_1a</i> | 1504.862       | 1001.779 | 1004     | 1         | —        | —               | 1004            | §            |
| <i>460_1a</i> | 1663.952       | 1040.196 | 647      | 1         | —        | —               | 647             | §            |

§—no solution after 3600s

**Table 4** Results of Exp 4

| Instance      | Initial iteration ( $k = 0$ ) |          | Algorithm                |          |          |          | Run time (s) |
|---------------|-------------------------------|----------|--------------------------|----------|----------|----------|--------------|
|               | $UB_F$                        | $LB_F$   | # iter.= $\# \Upsilon^k$ | $UB_F$   | $LB_F$   | $F^*$    |              |
| <i>10_1b</i>  | 6613.000                      | 6175.333 | 1                        | 6175.333 | 6175.333 | 6175.333 | 0.28         |
| <i>60_1b</i>  | 6036.720                      | 5297.495 | 1                        | 5297.495 | 5297.495 | 5297.495 | 0.59         |
| <i>110_1b</i> | 5271.851                      | 4385.000 | 1                        | 4385.000 | 4385.000 | 4385.000 | 0.62         |
| <i>160_1b</i> | 5871.734                      | 3481.623 | 1                        | 3481.623 | 3481.623 | 3481.623 | 2.83         |
| <i>210_1b</i> | 5087.836                      | 3451.424 | 1                        | 3451.424 | 3451.424 | 3451.424 | 2.38         |
| <i>260_1b</i> | 5327.763                      | 4031.825 | 1                        | 4031.825 | 4031.825 | 4031.825 | 242.60       |
| <i>310_1b</i> | 4810.362                      | 3424.941 | 1                        | 3424.941 | 3424.941 | 3424.941 | 8.16         |
| <i>360_1b</i> | 5038.987                      | 3457.085 | 1                        | —        | —        | —        | §            |
| <i>410_1b</i> | 4997.388                      | 3199.961 | 1                        | 3199.961 | 3199.961 | 3199.961 | 141.99       |
| <i>460_1b</i> | 5019.918                      | 3183.537 | 1                        | 3183.537 | 3183.537 | 3183.537 | 74.17        |

§—no solution after 1800s

being this  $LB_F$  equal to the  $F$ -optimal value ( $F^*$ ), the number of follower's feasible solutions computed ( $\# \Upsilon^k$ ) and the total run time (in seconds).

Table 2 shows the results of Exp. 2, differing from Exp.1 just in disabling the construction of the initial sample. The  $UB_F$  in iteration  $k = 0$  is the same as in Table 1, but the  $LB_F$  is, in general, different because only one solution results from iteration  $k = 0$  that can provide the lower-bound of  $F$ , which is  $(x^0, y_L^0)$ . In this experiment, only one solution is generated as input for iteration  $k = 1$ . Note that  $LB_F = -\infty$  for the instance *110\_1* at  $k = 0$ ; this may happen when  $(x^0, y_L^0)$  does not satisfy the upper-level constraints, i.e., it is bilevel infeasible. The algorithm could not compute the optimal solution for the biggest problem with 460 constraints within 1800 s. In this instance and some others, the construction of the initial sample reveals to be interesting for the good performance of the algorithm, which is still reinforced when the optimal solution is already included in the initial sample.

The results in these tables indicate that the algorithm finds the optimal solution in a reasonable computation time when the feasible region is not large.

The performance worsens significantly when the original feasible region is enlarged—Exp. 3 with results in Table 3. A very high number of solutions in the initial sample is computed, but disabling the utilization of the sample, in this case, does not provide results even with a significantly larger computation time. For the 10 test problems, the optimal solution was obtained just for the smallest one (10 variables and 4 constraints in each level) after 6 iterations, and the algorithm was able to generate bounds for the problem with 310 variables. For all other problems, the algorithm was not able to complete the iteration  $k = 1$  and did not provide solutions after 1 hour of computation.

In the special class of MILBLO problems (8), the algorithm displays a better performance (Table 4). The structure of the problem (in which lower-level constraints do not include upper-level variables) allows to simplify the sub-problems solved by the algorithm which do not require additional binary variables, as discussed in Sect. 3.1. The optimal solution was obtained after just one iteration, even without the initial sample, for all problems of Exp. 4 except instance *360\_1b* in which the algorithm could not complete the iteration  $k = 1$  within the computation budget.

## 5 Conclusions

In this paper, we analyzed the Lozano-Smith algorithm, which is a state-of-the-art approach based on an optimal-value-function reformulation, to solve the mixed-integer linear bilevel optimization problem. This algorithm solves a relaxation of the bilevel optimization problem with disjunctive constraints generated from optimal follower's responses to obtain upper bounds for the maximizing upper-level objective function. The lower bounds are derived from bilevel feasible solutions. Since all the subproblems solved by the algorithm are mixed-integer linear programs, off-the-shelf MIP solvers can be used without the need to implement more complex procedures. We show how the algorithm can be made more efficient when upper-level variables do not appear in the lower-level constraints.

We designed a set of computational experiments to assess the algorithm performance with different features in data instances of different dimensions and characteristics. The performance of the algorithm in the different groups of problems is discussed showing that it performs well for problems with moderate-size feasible regions, but the computation effort becomes impracticable when the feasible region expands. Nevertheless, the Lozano-Smith algorithm is very ingenious from theoretical and methodological perspectives, also displaying good results in practice for problems with particular characteristics.

**Acknowledgements** This work has been funded by FCT—Portuguese Foundation for Science and Technology within Projects UIDB/05037/2020 and UIDB/00308/2020.

## References

1. Moore, J.T., Bard, J.F.: The mixed integer linear bilevel programming problem. *Oper. Res.* **38**(5), 911–921 (1990)
2. DeNegre, S.T., Ralphs, T.K.: A branch-and-cut algorithm for Integer bilevel linear programs. In: Chinneck, J.W., Kristjansson, B., Saltzman, M.J. (eds.) *Operations Research and Cyber-Infrastructure. Operations Research/Computer Science Interfaces*, vol. 47, pp. 65–78. Springer, Boston (2009)
3. Xu, P., Wang, L.: An exact algorithm for the bilevel mixed integer linear programming problem under three simplifying assumptions. *Comput. Oper. Res.* **41**, 309–318 (2014)
4. Caramia, M., Mari, R.: Enhanced exact algorithms for discrete bilevel linear problems. *Optim. Lett.* **9**(7), 1447–1468 (2015)
5. Fischetti, M., Ljubić, I., Monaci, M., Sinnl, M.: A new general-purpose algorithm for mixed-integer bilevel linear programs. *Oper. Res.* **65**(6), 1615–1637 (2017)
6. Fischetti, M., Ljubić, I., Monaci, M., Sinnl, M.: Intersection cuts for bilevel optimization. In: Louveaux, Q., Skutella, M. (eds.), *Integer Programming and Combinatorial Optimization. IPCO 2016. LNCS*, vol. 9682, pp. 77–88. Springer, Cham (2016)
7. Zeng, B., An, Y.: Solving bilevel mixed integer program by reformulations and decomposition. *Optimization online*, 1–34 (2014). <https://optimization-online.org/wp-content/uploads/2014/07/4455.pdf>
8. Bolusani, S., Ralphs, T.K.: A framework for generalized Benders’ decomposition and its application to multilevel optimization. *Math. Program.* **196**(1–2), 389–426 (2022)
9. Mitsos, A., Lemonidis, P., Barton, P.I.: Global solution of bilevel programs with a nonconvex inner program. *J. Global Optim.* **42**(4), 475–513 (2008)
10. Lozano, L., Smith, J.C.: A value-function-based exact approach for the bilevel mixed-integer programming problem. *Oper. Res.* **65**(3), 768–786 (2017)
11. Soares, I., Alves, M.J., Henggeler Antunes, C.: A deterministic bounding procedure for the global optimization of a bi-level mixed-integer problem. *Eur. J. Oper. Res.* **291**(1), 52–66 (2021)
12. Chen, S., Feng, S., Guo, Z., Yang, Z.: Trilevel optimization model for competitive pricing of electric vehicle charging station considering distribution locational marginal price. *IEEE Trans. Smart Grid* **13**(6), 4716–4729 (2022)
13. Tavaslıoğlu, O., Prokopyev, O.A., Schaefer, A.J.: Solving stochastic and bilevel mixed-integer programs via a generalized value function. *Oper. Res.* **67**(6), 1659–1677 (2019)
14. Kleinert, T., Labbé, M., Ljubić, I., Schmidt, M.: A survey on mixed-integer programming techniques in bilevel optimization. *EURO J. Comput. Optim.* **9**, 100007 (2021)
15. Bard, J.F., Falk, J.E.: An explicit solution to the multi-level programming problem. *Comput. Oper. Res.* **9**(1), 77–100 (1982)
16. Bard, J.F.: *Practical Bilevel Optimization: Algorithms and Applications*. Kluwer Academic Publishers, Dordrecht (1998)
17. Mersha, A.G., Dempe, S.: Linear bilevel programming with upper level constraints depending on the lower level solution. *Appl. Math. Comput.* **180**(1), 247–254 (2006)
18. Vicente, L., Savard, G., Judice, J.: Discrete linear bilevel programming problem. *J. Optim. Theory Appl.* **89**, 597–614 (1996)

# Using Supplier Networks to Handle Returns in Online Marketplaces



Catarina Pinto, Gonalo Figueira, and Pedro Amorim

**Abstract** To encourage customers to take a chance in finding the right product, retailers and marketplaces implement benevolent return policies that allow users to return items for free without a specific reason. These policies contribute to a high rate of returns, which result in high shipping costs for the retailer and a high environmental toll on the planet. This paper shows that these negative impacts can be significantly minimized if inventory is exchanged within the supplier network of marketplaces upon a return. We compare the performance of this proposal to the standard policy where items are always sent to the original supplier. Our results show that our proposal—returning to a closer supplier and using a predictive heuristic for fulfilment—can achieve a 16% cost reduction compared to the standard—returning to the original supplier and using a myopic rule for fulfilment.

**Keywords** Returns management · Online marketplaces · Fulfillment

## 1 Introduction

Marketplace business models appear simple at first glance. Marketplaces facilitate transactions of items between suppliers and customers, retaining a margin in each sale. Suppliers are the owners of inventory and, therefore, responsible for sourcing and sharing their items in the platform [3]. We refer to marketplaces and platforms interchangeably in the paper. Once an order arrives, the marketplace confirms which suppliers can source the item and evaluate them according to different criteria—e.g. cost of shipping, distance to the customer, and other performance metrics that describe the supplier. The item is sent from the selected supplier to the customer and, in the case of a return, sent back to the former.

---

C. Pinto · G. Figueira · P. Amorim (✉)

INESC TEC, Faculdade de Engenharia da Universidade do Porto, R. Dr. Roberto Frias, 4200-465 Porto, Portugal

e-mail: [pamorim@fe.up.pt](mailto:pamorim@fe.up.pt)

Understandably, the location of suppliers and customers plays a significant role in marketplace logistics, both in terms of forward-shipping and return costs, and customer experience (e.g., lead time of the delivery). Location is especially important in global marketplaces with scarce items (i.e., seldom sourced by different suppliers). For example, it is not uncommon for a European customer of eBay to receive electronics from a supplier in China as the item was specific for this supplier. Industries such as fashion see their logistics problems exacerbated due to the high level of returns. As clothing fit and features are not easy to perceive without trying on the garments, the return levels in the industry hits 15%-40%, and is prone to grow as customers get more used to purchase online.<sup>1</sup> When an item is returned in a marketplace, this translates into 2x the number of trips (and costs) per item, which can often include transcontinental flights.

The financial and environmental impact of online returns does not come as news for retailers. For example, in 2018, the company Revolve recorded 499M\$ in sales while spending 531M\$ in returns after accounting for processing costs and lost sales [13]. Other companies, such as Amazon, Walmart, and Chewy, requested customers to keep the products they wished to return at zero cost during the 2020 Christmas holidays. For these players, some returns were more expensive to process and ship than the actual products, averaging 10\$ to 20\$ per return<sup>2</sup> [17]. Processing returns is often such a burden that retailers end up throwing away 25% of returned items (translating into over 2.3 billion kg of goods in landfills per year), a number expected to double in the next years [5]. As for the environment, these returns result in approximately 15 million tonnes of carbon emissions in transportation per year [13].

Understanding the challenges of global marketplaces with high return levels and scarce items, we discuss in this paper how to leverage the network of suppliers to reduce the financial and environmental impact of returns. We start by analyzing the choice of the sourcing supplier, comparing two selection approaches: one that selects the supplier based solely on current costs and one that considers the future demand and the probability of customer returns from different geographic areas. We name these policies *myopic* and *predictive*, respectively. Then, we question one critical assumption in the process. Namely, we ask, What if the item does not have to return to the sourcing supplier? What if the item can be sent to a supplier closer to the customer? In such a case, the distance traveled in the return leg might drastically decrease, reducing the impact on the company and the planet.

We base our experiments on a case study in luxury fashion and conduct our analysis for different scenarios of global dispersion of suppliers and consumers. We show that shipping costs can be greatly minimized if the platform can facilitate inventory exchanges within its supplier network upon a return. This can keep the inventory in the regions where demand is higher, decrease shipping costs, and improve customer experience. The cost reduction is even larger if the retailer adopts a predictive

---

<sup>1</sup> Due to the higher level of online purchasing during the COVID-19 pandemic, the number of returns in e-commerce jumped 70% in 2020 relative to 2019 [17].

<sup>2</sup> Excluding transportation.

approach that considers the expected future shipping costs in the fulfillment decision, rather than a myopic heuristic that assigns suppliers to customers only looking at the immediate costs. We test our idea in different supplier network designs and find that cost savings increase with supplier geographical dispersion. We further analyzed the sustainability impact of our proposal. We show that using the supplier network can save almost 12K tonnes of  $CO_2$  emissions in our use case.

Our contribution is twofold. First, we show that marketplaces are not using an important asset—their supplier network. We reveal that using this network for returns management significantly impacts financial and environmental costs for marketplaces. Second, we test our proposal in different supplier network designs, showing it performs best when suppliers are dispersed and closer to high-demand locations. This result is important for supply chain managers responsible for sourcing the necessary suppliers.

The rest of the paper is organized as follows. In Sect. 2, we review the literature related to our research. In Sect. 3, we describe the returns optimization problem of marketplaces. In Sect. 4 we explain our dataset and the different supply chain scenarios. In Sect. 5 we explore the impact of our policies on return costs and the environment, extending to the case where there is a warehouse in the marketplace supply chain. We conclude our paper in Sect. 6.

## 2 Literature Review

Online returns are a plague for many retailers. Despite their negative financial impact, retailers still see them as a necessary evil to drive demand to the online channel. Indeed, companies today must provide customers with the option to return (commonly) for free, which leads to even more returns.

Benevolent returns are so common that they can be seen as part of the customer purchasing process. For example, [12, 15] look at returns as part of the purchase episode between the retailer and the customer who aims to fulfill a need. This episode may result in zero, one, or multiple returns until the customer finds the right product or gives up. To minimize the cost of returns for the company, the authors use price and assortment to guide the retailer on how to handle return costs. Other authors propose to handle returns by directly reducing their frequency. For example, [9] use supply chain coordination contracts, giving retailers incentives to actively prevent (false) returns.

To minimize the number of returns, retailers started to actively dampen factors that foster online returns, such as the information available [6], existing customer reviews [18], and the return period granted by the retailer [7]. Instead of attempting to decrease the number of returns, this paper focuses on optimizing the fulfillment process.

Several authors have studied fulfillment optimization for the forward supply chain. For instance, [19] reevaluate the fulfillment plan of the not-yet-picked orders and rematches supply and demand to minimize costs, [11] aggregate orders before

allocating them to facilities, which is accomplished via a quasi-dynamic allocation policy, and [1] allocate orders as soon as they arrive while looking into the future by approximating expected future shipping costs. However, the former authors do not consider that products will be returned. In our paper, we use a heuristic inspired by the work of [1] and extend it to consider returns.

While reducing the financial impact of returns is important, its impact on sustainability must also be considered. In a 2018 study, [4] showed that, under most circumstances, online shopping is better for the planet than in-store shopping. However, excessive packaging [8], customer behaviour [14] that leads to higher delivery frequency, and product returns [4] negatively contribute to its sustainability. While practices such as waste management, recycling, or reuse have been standards for improving the impact of reverse logistics [14], in this paper we minimize the effect of the carbon footprint by decreasing the number of kilometers traveled in the reverse logistics process. Namely, we propose a new model where an internal market among suppliers is created, which enables items to be returned to a closer supplier (rather than the original one), improving the financial and environmental impact of returns.

### 3 Problem Formulation and Approach

Consider a global marketplace where both suppliers and customers are dispersed worldwide. When a customer orders, the platform identifies all the suppliers with the requested product  $n$  in stock. To reduce the granularity of our data, we group suppliers and customers by geography, representing the latter groups by  $i$  and  $j$ , respectively. The inventory available of  $n$  in a given geographical region depends on the number of suppliers in that region carrying  $n$  and the number of units carried by each supplier, and thus described by  $X_{i,n}$ .

After identifying the pool of suppliers with  $n$ , the platform can use a panoply of policies to select the sourcing supplier. A simple inventory allocation policy might only account for the forward shipping costs of the order, for example. In this case, the platform will select the supplier where  $c_{i,j}$ , the cost of shipping an item from a store in region  $i$  to a customer in region  $j$ , is minimum. We denominate such a policy by *Myopic Shipping*. As an item in inventory that is sold cannot be used to fulfill a future purchase, it might also be interesting to evaluate the future forward shipping costs, as did [1]. Imagine the case where two suppliers, one in Zaragoza, Spain and another in Bilbao, Spain, carry one unit of an item that a customer in Madrid, Spain, ordered. Zaragoza is closer to Madrid, so the platform selects this supplier to fulfill the order. However, a few minutes later, another order of the same item arrives from Barcelona, Spain. The item must now be sent from Bilbao. The difference between this myopic policy and a policy which would take into account the future demand

from Barcelona is more than 200 km.<sup>3</sup> We denominate such policy as *Predictive Shipping*.

The policies mentioned above (*Myopic Shipping* and *Predictive Shipping*) are focused on forward shipping costs. Analyzing the reverse flow, whenever a customer returns an item, the usual procedure is to send it to its original supplier. However, there might be customers nearby who will order the item shortly again. Instead of returning the item to the original supplier to then send it back again for a future purchase, we propose to leverage the network of the platform to attend future demand by returning items to suppliers which are *closer*<sup>4</sup> to the customer (and therefore cheaper) instead of the original supplier.

We formulate and compare, using different indicators, the four combinations of the above policies—Myopic Shipping with Return to Original Store ( $M_o$ ), Predictive Shipping with Return to Original Store ( $P_o$ ), Myopic Shipping with Return to Closer Store ( $M_c$ ), Predictive Shipping with Return to Closer Store ( $P_c$ ). Let us define each component of the policies individually.

**Myopic shipping.** In this policy, the marketplace uses only  $c_{ij}$  to decide on the best supplier  $i$ . The supplier region selected with the myopic policy for a given customer order can be defined by  $\arg \min_{i \in I} (c_{ij})$ , where  $I$  includes all regions with non negative inventory at the moment of the order.

**Predictive shipping.** Our predictive shipping policy extends the work of [1] to include product returns. The authors have already demonstrated that it is possible to improve forward shipping costs by accounting for future expected delivery costs.

The authors start by describing the structure of the optimal cost-to-go function  $J_n(X_n, t)$ , which represents the future expected cost to fulfill product  $n$  when its inventory position is  $X_n = [X_{in}]$ ,  $i \in \{1, \dots, I\}$  at moment  $t$ . The goal is to find the supplier  $i$  to source the item  $n$ , which minimizes the current shipping costs and the future expected shipping costs. To get a tractable solution, the problem is approximated by a linear transportation program, which minimizes the costs of shipping an item subject to its availability and demand. We expand the linear transportation program also to include product returns. Let the demand of product  $n$  over the current period  $\tau$  be  $\alpha_{jn}d_n\tau$ , where  $d_n$  is the average demand per day of product  $n$ ,  $\tau \in [t, t + 2T]$  (in days) is a moving window with double the size of the return period  $T$ , and  $\alpha_{j,n}$  is the ratio of the demand from region  $j$  over the demand of all regions, with  $\sum_j \alpha_{j,n} = 1$ ,  $\forall n$ . The probability of an item being returned is encoded in  $\rho_{j,n}(\tau) \in [0, 1]$ , which can be estimated by the ratio of the orders from product  $n$  returned in region  $j$  over the total number of orders from product  $n$  in the same

<sup>3</sup> Using Maps from Google, we estimate the distances between the cities to be Bilbao-Madrid: 403 km, Zaragoza-Madrid: 322 km, Bilbao-Barcelona: 608 km, Zaragoza-Barcelona: 306 km. Choosing to source the Madrid customer from Zaragoza leads to 930 km travelled while sourcing it from Bilbao leads to 709 km travelled, a 221 km difference.

<sup>4</sup> Due to deals with third-party logistic companies, the closest store is not always the less expensive to ship to. Throughout this paper, we use the term *closer* to indicate that stores with higher proximity to the customer are usually less expensive than stores further away.

region. Finally, we use  $w_{ij}$  as the decision variable of the linear transportation program, which represents the flow between supplier region  $i$  and customer region  $j$ . The resulting linear transportation program reads as follows (note that  $n$  is dropped for simplicity and the cost matrix is symmetric, i.e.,  $c_{ij} = c_{ji}$ ):

$$\min_w \sum_{i,j} c_{ij} w_{ij} + \sum_{i,j} \rho_j c_{ij} w_{ij} \quad (1)$$

subject to

$$\sum_j w_{ij} \leq X_i \quad \forall i \quad (2)$$

$$\sum_i w_{ij} = \alpha_j d\tau \quad \forall j \quad (3)$$

$$w_{ij} \geq 0 \quad \forall i, j$$

The first constraint, Eq. (2), ensures that the amount of supply leaving a supplier region is less or equal to the amount available in that region. The second constraint, Eq. (3), ensures that demand in a given customer region is met. The formulation assumes that the available supply is sufficient to meet the demand in the moving window  $\tau$ . Otherwise, we change the expected demand in that period,  $d$ , to match the available inventory level. The cost-to-go function is then approximated by the dual value,  $s_i(X)$ , associated with the inventory constraint of the specific item at supplier region  $i$  (i.e., Eq. (2)). The best supplier is the one that minimizes the sum of the current and future shipping costs, i.e.,  $\arg \min_{i \in I} (c_{ij} - s_i(X))$ . Note that the dual variable is non-positive, and hence when subtracted is either increasing the total cost or maintaining the same value.

**Return to original store.** As mentioned, it is common practice to send returned products back to their original supplier—this policy. However, such a policy means that if a product is sent from London to Sydney, it must return to London, even if local customer orders are coming soon for the same product. Model (1)–(3) incorporates this traditional policy.

**Return to a closer store.** This policy returns products to the cheapest store, minimizing return shipping costs. This store is usually close to the customer. In order to convert model (1)–(3) to this new situation, we simply replace  $c_{ij}$  by  $c_{L_j,j}$  in the second term of objective function (1), where  $L_j$  is the cheapest (*closer*) store region to customer region  $j$ . When using the Predictive Shipping model the optimization function is therefore rewritten as:

$$\min_w \sum_{i,j} c_{ij} w_{ij} + \sum_{i,j} \rho_j c_{L_j,j} w_{ij} \quad (4)$$

## 4 Framework for Evaluation

This section discusses the data provided to us and the supplier network structure scenarios used throughout the paper.

### 4.1 Data Overview

In this work, we partnered with a multi-brand luxury fashion retailer that runs an online marketplace. The company does not own stock, working in a drop-shipping model and assuring that products go from their original supplier to the customer that requested them using a third-party logistics partner. The data provided by our partner includes the purchase history between 2016-06-01 and 2017-12-31. The available data can be distinguished into order-related data (e.g., order value and items), customer-related data (e.g., customer country and city), and product-related data (e.g., product type and category). To evaluate the impact of the different fulfilment and return policies for different return levels, we selected 120 products with varying percentages of return (corresponding to the ratio between the number of times the item was returned by the number of times the item was ordered), with a return percentage *evenly* distributed between 2% and 25%. These products were mainly clothes (43.3%), shoes (47.5%), lifestyle items, bags, and accessories. All products were from big brands like Kenzo, Valentino, Chiara Ferragni, Saint Laurent, Gucci, or Adidas, with prices ranging from 20€ (belt) to 800€ (denim). More than half of the returns originated from 12 regions,<sup>5</sup> including California (11.29%), Australia (10.68%), United Kingdom, New York, Germany, Russian Federation, Brazil, China, France, Hong Kong, Korea, and the Netherlands. To compare the performance of the different policies across products, we selected the first 150 orders of each item and assumed a similar amount of inventory in the network. The orders and the respective customers are kept constant across the different experiments.

We rely on three assumptions to evaluate the different policies. First, the inventory available in each supplier is known to the platform [2]. Second, the platform can forecast demand and return rates with high accuracy. This assumption is reasonable given a large amount of supply and demand data from marketplaces. Third, as most of our partner's returns are not due to product damage, we assume that items become available for sale after a return (we use a lead time  $T$  of 5 days based on conversations with the firm).

---

<sup>5</sup> We consider a region to correspond to a country or a United States state.

## 4.2 Supplier Network Design

Platforms need to balance the two sides of the market—customers and suppliers—to grow. Deciding where to recruit suppliers depends on the type of industry (for example, luxury fashion producers are more concentrated in Europe), the location of the customers, the platform, and the interest of the suppliers to participate. The supplier network is flexible as these factors vary with time and the company’s growth stage. Therefore, to make our insights applicable to distinct marketplaces, we decided to detach from our partner network (which has also changed substantially since 2016). We use three scenarios with different locations, as shown in Fig. 1. In the first one, all suppliers are in a single European hub. In the second, we add a few suppliers on the United States East Coast. Lastly, we add three extra suppliers in high-demand locations, such as Hong Kong, Russia, and California.

As inventory management in luxury fashion fits a newsvendor model (suppliers have one shot at selling their seasonal inventory with little to no replenishment during a season), we used a single supply period without opportunity for replenishment. The initial inventory level is fixed at 150 units for all the scenarios. This assumption is relaxed in Sect. 5.2.



**Fig. 1** Supply Network Design Scenarios. The size of the blue bubbles corresponds to the number of units of inventory available in the different locations. The total inventory is set to 150 units in the three cases

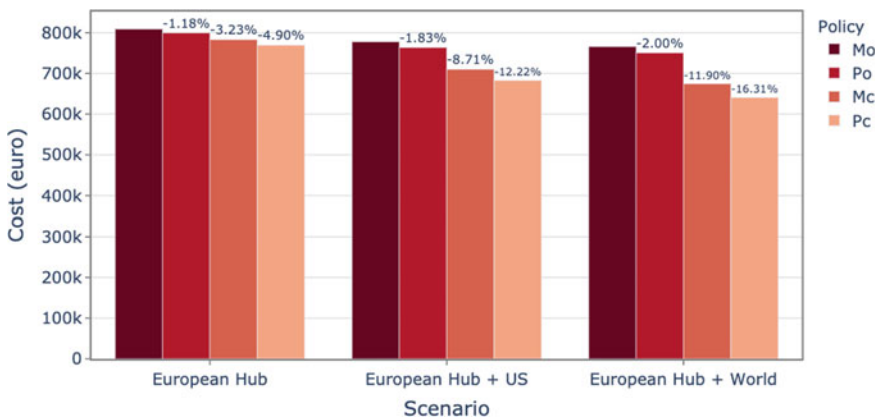
## 5 Leveraging the Supplier Network

In this section, we demonstrate the effectiveness of our proposal—Predictive Shipping with Return to Closer Store—in reducing shipping and environmental costs. We further perform a sensitivity analysis of the different policies over distinct inventory levels.

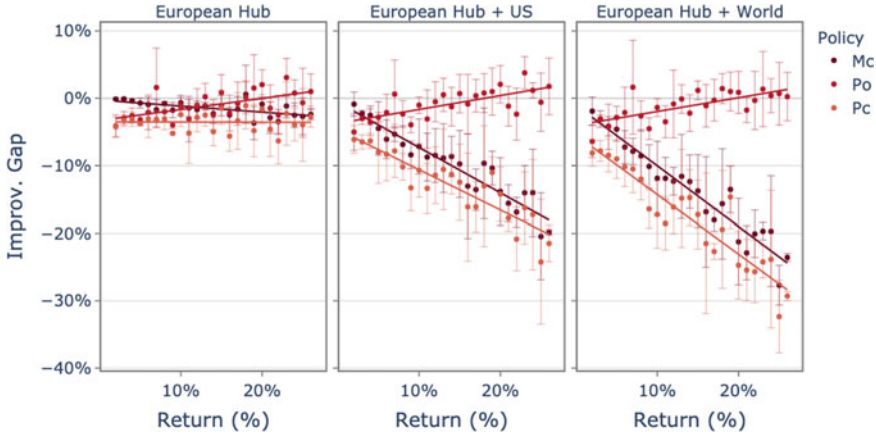
### 5.1 Policies Comparison

We start by evaluating the performance of the four fulfillment policies  $M_o$ ,  $M_c$ ,  $P_o$ , and  $P_c$ , for the three configurations described above (Sect. 4.2). Our results are depicted in Fig. 2. The percentage value on the top of each bar corresponds to the *improvement gap* w.r.t. the  $M_o$  policy from the same scenario, and is calculated by  $(C^E - C^{M_o})/C^{M_o}$ , where  $E$  corresponds to the evaluated policy. By analyzing the chart, we can see that there is a reduction in costs as the dispersion of the supply network increases. For the policies that return to the original store, we verify that the  $P_o$  records slightly better results than  $M_o$ , which is in line with the values presented by [1], who reported a reduction in outbound shipping costs in the order of 1%. From a different perspective, the policies that route returned inventory to closer suppliers outperform those which return to the original supplier. Specifically, they reduce shipping costs between 3.23% and 11.90% for  $M_c$ , and between 4.90% and 16.31% for  $P_c$ .

We further analyze the performance evolution of the different policies by the return level (c.f. Fig. 3). We again use the *improvement gap* as described before. We conjecture that larger return rates correspond to higher savings, on average.



**Fig. 2** Total shipping costs obtained with the proposed policies for the different supply network scenarios



**Fig. 3** Improvement gap of shipping costs by return level for the different supplier network scenarios

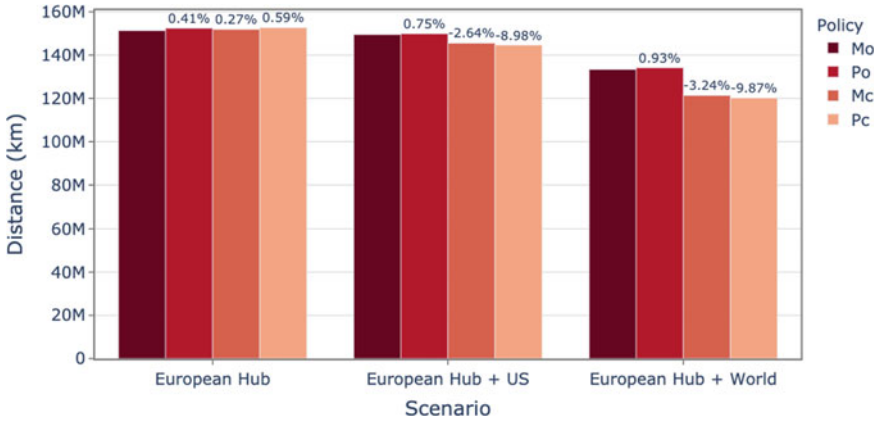
Let us start by analyzing the different scenarios. When there is only one hub, the performance of the different policies goes from almost no improvement to a reduction of 5% in costs. However, the difference in performance between the different policies is not so evident. In the scenarios where inventory is more dispersed, the cost savings with  $M_c$  and  $P_c$  substantially increase with an increased return level. Specifically,  $P_c$  achieves reductions of 17% and 23% for *European Hub + US* and *European Hub + World* scenarios, respectively. Overall, the results for the different return levels reiterate the interest of  $P_c$  across the different possible networks. The advantage of this policy grows with higher return rates. On the other hand, the performance of current policies, such as  $P_o$ , deteriorates with increasing return rates.

## 5.2 Sensitivity Analysis

So far, we have assumed that supply levels equal the demand levels observed throughout the season. However, in practice, this is often not the case. Therefore, we tested the sensitivity of our policies in scenarios when there is 20% less and 20% more inventory at the season start. Our results are presented in Table 1. We show that the performance of our policies is reduced for the single European Hub when there is excess inventory but amplified in the case of shortages. In the other scenarios, our policies perform similarly or better than  $M_o$ . Therefore, we conclude that our proposed policy,  $P_c$ , performs well in excess or scarce inventory cases, especially when suppliers are geographically dispersed.

**Table 1** Sensitivity analysis for different levels of inventory. Both total shipping costs and improvement gaps are reported

| Inventory Level | European Hub |         |         |         | European hub and US stores |         |         |         | European hub and world stores |         |         |         |
|-----------------|--------------|---------|---------|---------|----------------------------|---------|---------|---------|-------------------------------|---------|---------|---------|
|                 | $M_o$        | $M_c$   | $P_o$   | $P_c$   | $M_o$                      | $M_c$   | $P_o$   | $P_c$   | $M_o$                         | $M_c$   | $P_o$   | $P_c$   |
| -20% units      | 815,913      | 798,084 | 786,680 | 767,218 | 758,228                    | 674,583 | 722,747 | 629,974 | 687,694                       | 585,551 | 654,027 | 538,580 |
|                 |              | -2.19%  | -3.58%  | -5.97%  |                            | -11.03% | -4.68%  | -16.91% |                               | -14.85% | -4.90%  | -21.68% |
| 150 units       | 858,388      | 845,431 | 850,911 | 828,541 | 786,100                    | 707,797 | 780,211 | 684,664 | 698,713                       | 603,238 | 692,244 | 574,450 |
|                 |              | -1.51%  | -0.87%  | -3.48%  |                            | -9.96%  | -0.75%  | -12.90% |                               | -13.66% | -0.93%  | -17.78% |
| +20% units      | 835,303      | 827,843 | 852,433 | 828,671 | 750,619                    | 679,779 | 768,156 | 676,875 | 647,907                       | 562,313 | 661,988 | 551,812 |
|                 |              | 1.46%   | 4.48%   | 1.56%   |                            | -10.35% | 1.31%   | -10.73% |                               | -18.23% | -3.74%  | -19.76% |



**Fig. 4** Total distance travelled in km for the different supply network scenarios

### 5.3 Impact on Sustainability

One of the primary focuses of every company today is to reduce carbon emissions. Transportation is a big part of the problem, as it is challenging to minimize transportation costs when customers request faster and more frequent shipping [16]. We also analyze our policies' impact on the number of kilometers (km) traveled. Figure 4 shows that, as expected, the total number of km decreases as suppliers are more dispersed using the proposed policies. As mentioned before, distance does not hold a 1:1 relationship with costs. As we use historical costs that account for special deals made between our partner and the third-party logistic companies, frequently, the regions implying lower costs are not the closest ones. This particularity is the reason behind the variation observed in the *European Hub* scenario in Fig. 4.

In 2018, the DPDgroup, the largest parcel delivery network in Europe, estimated an average of 0.83kg of CO<sub>2</sub> emitted per parcel delivered [10]. Notice that fashion marketplaces, parcel delivery is the usual mode of transportation. Using this value, we estimate that the *P<sub>c</sub>* policy saves more than 4K tonnes of CO<sub>2</sub> in the *European Hub+US* scenario and almost 12K tonnes of CO<sub>2</sub> in the *European Hub+World* scenario. This result strengthens the added value of pursuing this policy.

## 6 Conclusions

Supplier networks are an underrated asset of marketplaces. Indeed, we show in our study that these networks can be leveraged to reduce reverse shipping costs and transportation's carbon footprint. Specifically, we find that by routing returned products to closer suppliers (rather than the original supplier that provided the item), the platform can achieve cost savings in the order of 1–16%, depending on the policy

used. We also demonstrate the impact of these policies in reducing the distance travelled, which maps to the CO<sub>2</sub> emission, recording almost 10% reduction with the best performing policy.

Our insights have a clear impact on supplier network design in online marketplaces. We show that, when using a return to closer store policy, marketplaces may establish a strong supplier network in a given region (in our case study, in Europe, where most suppliers are). Still, they should also maintain a small number of suppliers dispersed in the rest of the world where supply is harder to find.

While showing that using the supplier network can yield several benefits to a marketplace's value chain, this research work also opens avenues for further research, namely, in understanding how to incentivize the closer suppliers to receive these items. Several different mechanisms can be tested to this end. For example, the platform could use the shipping savings to incentivize suppliers to receive products they have not planned for.

## 7 Appendix: Nomenclature for Mathematical Formulation

### Indices

- $i$  suppliers  
 $j$  customers

### Parameters

- $X_i$  inventory carried by supplier  $i$   
 $c_{i,j}$  cost of shipping an item from a store in region  $i$  to a customer in region  $j$   
 $c_{L_j,j}$  cost of shipping an item from a customer in region  $j$  to the cheapest (closer) store  
 $d_{\tau}$  average demand in the moving forecast window  
 $\alpha_j$  ratio of the demand from region  $j$  over the demand of all regions  
 $\rho_j$  probability of a good being returned in region  $j$

### Decisions Variables

- $w_{ij}$  flow of products between supplier region  $i$  and customer region  $j$

## References

1. Acimovic, J., Graves, S.C.: Making better fulfillment decisions on the fly in an online retail environment. *Manufact. Serv. Oper. Manag.* **17**(1), 34–51 (2015)
2. Martinez-de Albeniz, V., Pinto, C., Amorim, P.: Driving supply to marketplaces: optimal platform pricing when suppliers share inventory. Available at SSRN 3643261 (2020)
3. Martínez-de Albéniz, V., Pinto, C., Amorim, P.: Driving supply to marketplaces: optimal platform pricing when suppliers share inventory. *Manufact. Serv. Operat. Manag.* (2022)
4. Bertram, R.F., Chi, T.: A study of companies' business responses to fashion e-commerce's environmental impact. *Int. J. Fash. Design, Technol. Educ.* **11**(2), 254–264 (2018)
5. CNBC: That sweater you don't like is a trillion-dollar problem for retailers. these companies want to fix it. <https://www.cnbc.com/2019/01/10/growing-online-sales-means-more-returns-and-trash-for-landfills.html> (2019). Accessed: 2021-03-21
6. De, P., Hu, Y., Rahman, M.S.: Product-oriented web technologies and product returns: an exploratory study. *Inf. Syst. Res.* **24**(4), 998–1010 (2013)

7. Ertekin, N., Ketzenberg, M.E., Heim, G.R.: Assessing impacts of store and salesperson dimensions of retail service quality on consumer returns. *Prod. Oper. Manage.* **29**(5), 1232–1255 (2020)
8. Escursell, S., Llorach-Massana, P., Roncero, M.B.: Sustainability in e-commerce packaging: a review. *J. Cleaner Product.* **280**, 124,314 (2021)
9. Ferguson, M., Guide Jr, V.D.R., Souza, G.C.: Supply chain coordination for false failure returns. *Manufact. Serv. Operat. Manage.* **8**(4), 376–393 (2006)
10. Group, D.: DPD Group remains #1 on the voluntary carbon offset market in the CEP sector. <https://www.dpd.com/group/en/2019/06/04/dpdgroup-remains-1-on-the-voluntary-carbon-offset-market-in-the-cep-sector/> (2018). (Accessed on 05/26/2021)
11. Mahar, S., Wright, P.D.: The value of postponing online fulfillment decisions in multi-channel retail/e-tail organizations. *Comput. Operat. Res.* **36**(11), 3061–3072 (2009)
12. Samorani, M., Alptekinoglu, A., Messinger, P.R.: Product return episodes in retailing. *Serv. Sci.* **11**(4), 263–278 (2019)
13. Schiffer, J.: The unsustainable cost of free returns. <https://www.voguebusiness.com/consumers/returns-rising-costs-retail-environmental> (2019). Accessed: 2021-03-21
14. Van Loon, P., Deketele, L., Dewaele, J., McKinnon, A., Rutherford, C.: A comparative analysis of carbon emissions from online retailing of fast moving consumer goods. *J. Clean. Prod.* **106**, 478–486 (2015)
15. Wagner, L., Martínez-de Albéniz, V.: Pricing and assortment strategies with product exchanges. *Oper. Res.* **68**(2), 453–466 (2020)
16. Wagner, L., Pinto, C., Amorim, P.: On the value of subscription models for online grocery retail. *Eur. J. Operat. Res.* (2020)
17. Wall Street Journal: Amazon, walmart tell consumers to skip returns of unwanted items. <https://www.wsj.com/articles/amazon-walmart-tell-consumers-to-skip-returns-of-unwanted-items-11610274600> (2021). Accessed: 2021-03-21
18. Wang, Y., Ramachandran, V., Liu Sheng, O.R.: Do fit opinions matter? the impact of fit context on online product returns. *Inf. Syst. Res.* **32**(1), 268–289 (2021)
19. Xu, P.J., Allgor, R., Graves, S.C.: Benefits of reevaluating real-time order fulfillment decisions. *Manufact. Serv. Operat. Manage.* **11**(2), 340–355 (2009)

# The Multiple Ambulance Type Dispatching and Relocation Problem: Optimization Approaches



A. S. Carvalho and M. E. Captivo

**Abstract** Ambulance dispatching and relocation decisions in the Emergency Medical Services (EMS) context are crucial in minimizing emergency response times and, consequently, in avoiding dangerous consequences to the patients' health. Several ambulance types exist, which vary in the equipment and crews, and should be considered depending on the emergency severity. We propose optimization approaches to solve these problems that incorporate a preparedness measure which aim is to achieve a good service level for current and future emergencies. As a case study we consider the Portuguese EMS for the Lisbon area. The current strategy dispatches the closest ambulance and always relocates ambulances to home bases. We highlight the potential of using preparedness and allowing relocations to any base since we can improve the current system's response. To help in the decision-making process we propose a decision-support system which embeds the proposed approaches.

**Keywords** Emergency medical service · Ambulance management · Dispatching · Relocation

## 1 Introduction

Emergency Medical Service (EMS) is one of the most important healthcare services. It aims to serve emergencies with an effective response and has to manage and mobilize several resources. EMS provides basic medical care for any person in an emergency situation and corresponds to the pre hospital part.

Figure 1 shows the EMS chain of events. Dispatching decisions occur after the reporting when someone calls the emergency number. When an ambulance is needed (in some cases, medical advice by phone may be enough), one or more ambulances should be chosen from the set of available ambulances. We consider multiple ambulance types, each type has different specialized equipment and it is crewed by different

---

A. S. Carvalho (✉) · M. E. Captivo  
CMAFcIO, Faculdade de Ciências, Universidade de Lisboa, Lisbon, Portugal  
e-mail: [asofia.fcarvalho@gmail.com](mailto:asofia.fcarvalho@gmail.com)  
URL: <https://ciencias.ulisboa.pt/pt/cmafcio>

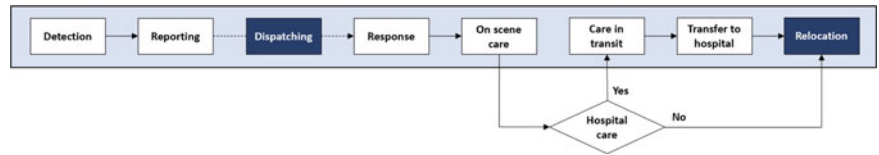


Fig. 1 EMS chain of events

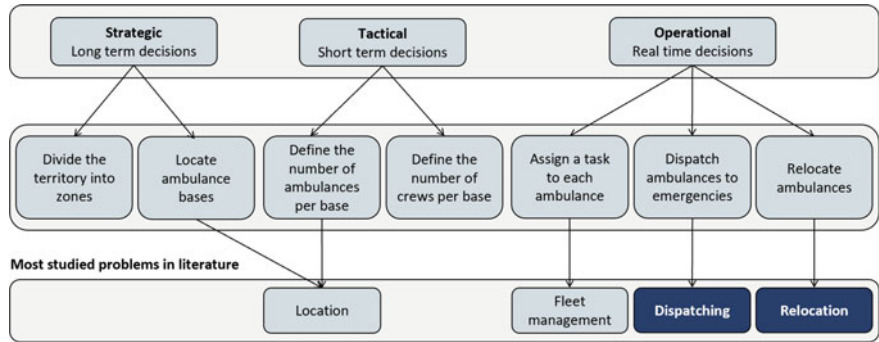


Fig. 2 EMS main problems

medical professionals. Depending on the emergency severity, the more appropriate ambulance type should be chosen. Relocation can occur in two different situations: (i) when the ambulance finishes the emergency service and should be relocated to a base or dispatched to an emergency that occurs meanwhile; and (ii) when an available ambulance changes its location to balance the ambulances’ distribution in the area under study. Dispatching and relocation tasks are the focus of this work and correspond to the operational level of decisions.

EMS is one of the most studied healthcare areas in which Operations Research (OR) is applied. These studies face different EMS problems in the strategic, tactical and operational level as presented in Fig. 2. Concerning the proposed approaches, the literature is mainly divided into exact methods where the first studies were proposed by [12, 34] as deterministic static models; heuristic algorithms (see [13, 19, 28, 32]); and simulation (see [20]).

The evolution in EMS studies has been considered in many review papers. Reference [31] developed one of the first reviews. Other review papers were presented by [1, 5, 7, 10, 14, 24, 27, 29, 30].

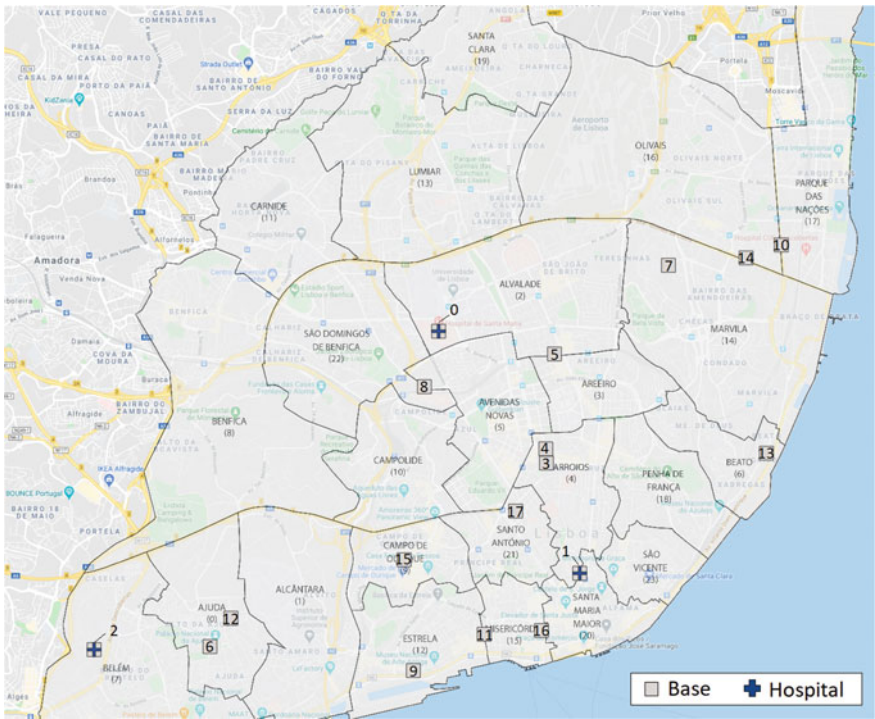
Literature concerning dispatching and relocation decisions is described with more detail in Sect. 3.

## 2 Contextualization: The Portuguese Case

As a case study, we consider the Portuguese EMS operation in the Lisbon area, where there are 18 bases where ambulances are parked and 3 reference hospitals (and simultaneously bases) as shown in Fig. 3. The 24 Lisbon districts are considered as sub-zones. Each sub-zone has an average number of expected ambulances per type known from historical data.

The current Portuguese EMS rules for dispatching and relocation in Lisbon area are as follows. The dispatching rule is to always dispatch the closest available ambulance of the required type. The relocation rule is to relocate ambulances to a fix base, the *home base*.

Concerning ambulances, each ambulance has an associated working shift. There are two types of ambulances: advanced life support (ALS) and basic life support (BLS) ambulances. Table 1 shows the existing bases and the ambulances with the corresponding base as home. Furthermore, for each ambulance, the corresponding working shift is presented.



**Fig. 3** Lisbon area divided into districts with the corresponding bases (identification number (id) is presented) and hospitals location

**Table 1** Available ambulances per base in Lisbon

| Base id | Number of ambulances |     | Working shifts                                                               |
|---------|----------------------|-----|------------------------------------------------------------------------------|
|         | ALS                  | BLS |                                                                              |
| 0       | 1                    | 0   | All day ( $\times 1$ )                                                       |
| 1       | 1                    | 0   | All day ( $\times 1$ )                                                       |
| 2       | 1                    | 0   | All day ( $\times 1$ )                                                       |
| 3       | 1                    | 0   | All day ( $\times 1$ )                                                       |
| 4       | 0                    | 2   | All day ( $\times 2$ )                                                       |
| 5       | 0                    | 3   | All day ( $\times 2$ ) / 8 AM–0 AM ( $\times 1$ )                            |
| 6       | 0                    | 2   | All day ( $\times 1$ ) / 8 AM–0 AM ( $\times 1$ )                            |
| 7       | 0                    | 3   | All day ( $\times 1$ ) / 8 AM–0 AM ( $\times 1$ ) / 4 PM–0 AM ( $\times 1$ ) |
| 8       | 0                    | 1   | 8 AM–4 PM ( $\times 1$ )                                                     |
| 9       | 0                    | 1   | All day ( $\times 1$ )                                                       |
| 10      | 0                    | 1   | 8 AM–0 AM ( $\times 1$ )                                                     |
| 11      | 0                    | 1   | All day ( $\times 1$ )                                                       |
| 12      | 0                    | 3   | All day ( $\times 3$ )                                                       |
| 13      | 0                    | 6   | All day ( $\times 6$ )                                                       |
| 14      | 0                    | 2   | All day ( $\times 2$ )                                                       |
| 15      | 0                    | 1   | All day ( $\times 1$ )                                                       |
| 16      | 0                    | 2   | All day ( $\times 2$ )                                                       |
| 17      | 0                    | 4   | All day ( $\times 4$ )                                                       |

Emergencies may have one of two types: type I and type II. Type I emergencies are the more severe ones and need, at least, 1ALS + 1BLS ambulances, and should be reached within 15 min. Type II emergencies need, at least, 1BLS ambulance and should be reached within 19 min. Emergencies may occur in any location of the area under study.

Concerning the Portuguese EMS, many patients argue that the response time is too high which has direct consequences in their health in some situations. Besides this, the importance of having an effective and efficient response is an issue that concerns society. Thus, we propose optimization approaches to solve the ambulance dispatching and relocation problem. Furthermore, dispatching and relocation decisions are currently a handmade task. EMS dispatchers identified the need to have a better system to dispatch ambulances to emergencies and to relocate them. To face this, we propose a decision-support tool using Geographic Information System which embeds the proposed approaches.

The Portuguese EMS case for the Lisbon area was previously studied in [11] where a simpler version of the real ambulance dispatching and relocation Portuguese EMS problem was considered: few real-life features were introduced and only a single ambulance type was considered. In the present work, we generalize the ambulance dispatching and relocation problem to face the real Portuguese EMS problem. This paper innovates by considering:

1. The existence of working shifts as it is a Portuguese EMS specific feature. Although referred to in [37], it is not common to refer to ambulances' working shifts in the literature.
2. A new *preparedness* measure which includes multiple ambulance types, i.e., we consider the demand for different ambulance types instead of emergencies' demand in which different needs of EMS resources are not identified.
3. The *system response time* measure, which makes the understanding of the preparedness impact in an EMS system more intuitive.
4. The existence of uncovered and partial covered emergencies from previous time periods. Uncovered emergencies are referred to in some works, but, in most cases, it is assumed that the existing ambulances are always enough to face the emergencies' occurrence. The possibility of having partially covered emergencies when only a subset of the ambulances needed was dispatched is not usually referred to in the literature.

### 3 Ambulance Dispatching and Relocation

At the operational level, EMS managers face the dispatching and relocation decisions. As explained before, these decisions integrate the EMS chain of events. Figure 4 relates the emergency occurrence with the ambulances' location and status. It also represents the duration of the underlying activities. The emergency *service time* is the duration of the emergency service between the dispatching and the relocation. The *response time* is the time between the ambulance dispatching and the moment the ambulance reaches the emergency. A *maximum response time* is recommended to avoid dangerous consequences in patients' health. Above this value, the patient's life may be at risk, and, if needed, we consider the *extra time* to reach the emergency location, which corresponds to the difference between the response time and the maximum response time.

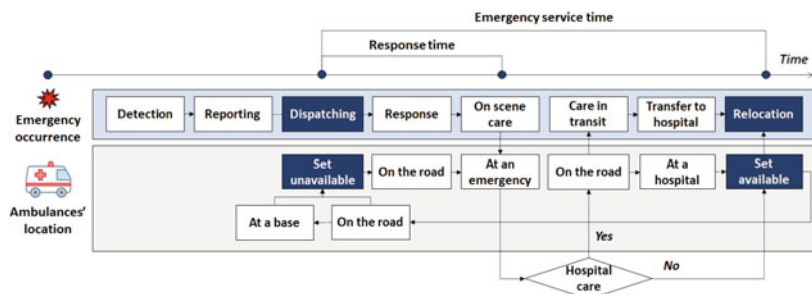


Fig. 4 Dispatching and relocation decisions in the EMS chain of events

In the literature, several works focus the dispatching and relocation decisions. Dispatching problem was studied by [20] who present a simulation-based analysis tool that includes location, dispatching, and fleet management decisions. Reference [8] present a recursive simulation-optimization framework to solve the ambulance location and dispatching problem. Reference [18] present a comparison between different dispatching policies, and [17] show that the closest policy, commonly used in the EMS systems, is quite far from optimal when minimizing the number of emergencies reached out of the maximum response time.

Concerning the relocation problem, [2, 3, 36] consider the use of *compliance tables*. [21] present one of the first models, specifically designed for New York City fire companies. Reference [25] use an approximate dynamic programming approach. Reference [39] propose a simulation-optimization method to determine the number of relocations needed to face the number of emergencies. Reference [35] evaluate the impact of the number of relocations in the service provider's performance. Reference [16] study the dynamic ambulance relocation problem with the goal of maximizing the number of emergencies that are reached within the maximum response time. Reference [37] highlight the impact of a dynamic relocation policy in reducing the extra times.

Many authors focus on both the dispatching and relocation problems (see [6, 9, 15, 26, 33]).

In this work, we study the *multiple ambulance type dispatching and relocation problem* (M-ADRP). It is solved under the following assumptions:

- An emergency needs a certain number of ambulances for each ambulance type. All the ambulances dispatched must respect the corresponding type.
- An ambulance can only be dispatched if the current time is within its working shift.
- After the dispatching decision, the ambulance cannot be reassigned to a different emergency before the current emergency service has finished.
- After on-scene care, if hospital care is needed, the patient is transported to the reference hospital of the sub-zone where the emergency occurred, which is a pre-defined and fixed hospital for each sub-zone.
- A base must be assigned to ambulances that have finished the emergency service, i.e., the ones located at an emergency or at a hospital.
- Available ambulances located at bases or on the road may change the current assigned base if the current time is within the corresponding working shift.
- An ambulance can change the current target base if  $\eta$  periods had passed from the last target base change for that ambulance.
- Available ambulances that are not in their working shifts (although the last decision was taken within the working shift) can only be relocated to their home base.

### 3.1 Notation

Dispatching and relocation decisions are solved at a given period  $t$ . The notation used is presented in Tables 2 and 3 which defines the sets and parameters, respectively.

### 3.2 System Coverage

To analyze system's coverage, we use a *preparedness* measure. According to [4], “preparedness is a way of evaluating the ability to serve potential patients with ambulances now and in the future”. It is a common measure in the literature to evaluate system's coverage (see [4, 11, 22, 23]). A preparedness measure considers a partition of the area under study into sub-zones, each one represented by a reference point in order to compute distances to/from the sub-zone. We extend the preparedness measure proposed by [23] to deal with different ambulance types. The *time-preparedness-ambulance-type* measure is defined for a time period  $t$  as:

$$p'_{AK} = \frac{1}{\sum_{q \in Z} \sum_{k \in K} \lambda_{qk} \times (1 + \min_{i \in A: \kappa_i = k} \{d_{iq}\})} \quad (1)$$

The preparedness measure denominator is defined as the *system response time*. The goal is to *minimize the system response time* on contrary to the system preparedness, which goal is to maximize.

**Table 2** Sets

|                                   |                                                                                                                                          |
|-----------------------------------|------------------------------------------------------------------------------------------------------------------------------------------|
| $A$                               | Available ambulances                                                                                                                     |
| $A^{WS}$                          | Available ambulances within the corresponding working shift                                                                              |
| $A^{\overline{WS}}$               | Available ambulances that are not within the corresponding working shift (although the last decision was taken within the working shift) |
| $A^B$                             | Available ambulances located at bases                                                                                                    |
| $A^R$                             | Available ambulances located on the road                                                                                                 |
| $A^{Target}$                      | Available ambulances that can change its target base                                                                                     |
| $E^{Partial\ Covered}$            | Emergencies with, at least, one ambulance previously assigned                                                                            |
| $E^{\overline{Partial\ Covered}}$ | Emergencies with no ambulance previously assigned                                                                                        |
| $E$                               | Emergencies to be analyzed, $E = E^{Partial\ Covered} \cup E^{\overline{Partial\ Covered}}$                                              |
| $B$                               | Bases                                                                                                                                    |
| $Z$                               | Sub-zones                                                                                                                                |
| $K$                               | Ambulance types                                                                                                                          |

**Table 3** Parameters

|                      |                                                                                                                                                                                                              |
|----------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $h_i$                | Home base of ambulance $i$ , $i \in A$                                                                                                                                                                       |
| $\eta_j^k$           | Number of ambulances of type $k$ needed in emergency $j$ , $k \in K, j \in E$                                                                                                                                |
| $\kappa_i$           | Type of ambulance $i$ , $i \in A$                                                                                                                                                                            |
| $d_{ab}$             | Travel time between location of $a$ and location of $b$ at $t$ , $a, b \in A \cup E \cup B \cup Z$                                                                                                           |
| $b_i$                | Target base for available ambulance $i$ at $t$ , $i \in A^B \cup A^R$                                                                                                                                        |
| $\lambda_{qk}$       | Number of ambulances of type $k$ needed in sub-zone $q$ at $t + 1$ , $k \in K, q \in Z$                                                                                                                      |
| $\sigma_j$           | Accumulated waiting time for emergency $j$ at $t$ , i.e. elapsed time since the emergency occurred, $j \in E$                                                                                                |
| $\delta$             | Average setup time needed before an ambulance starts to travel to an emergency                                                                                                                               |
| $R_j^{MAX}$          | Maximum response time for emergency $j$ , $j \in E$                                                                                                                                                          |
| $Target^{MAX}$       | Maximum number of ambulances that can change its target base at $t$                                                                                                                                          |
| $p_k^{Ambulance}$    | Penalization for each lacking ambulance of type $k$ for an emergency, i.e. the difference between the number of ambulances needed of each type and the actual number of ambulances assigned to it, $k \in K$ |
| $p^{Uncovered}$      | Penalization for each uncovered emergency                                                                                                                                                                    |
| $p^{Response}$       | Penalization for reaching an emergency $j$ with a response time higher than $R_j^{MAX}$ , $j \in E$                                                                                                          |
| $p^{Target}$         | Penalization for changing the current target base                                                                                                                                                            |
| $INF^{Travel\ time}$ | Represents a huge travel time. It is set to $2 \times \max\{d_{lq} : l \in B \wedge q \in Z\}$                                                                                                               |
| $M^{Travel\ time}$   | $\max\{d_{ab} : a, b \in A \cup E \cup B \cup Z\}$                                                                                                                                                           |

## 4 Approaches

We propose two optimization approaches: a mixed-integer linear programming model and a heuristic, described in Sects. 4.1 and 4.2, respectively.

### 4.1 Mixed-Integer Linear Programming Model

M-ADRP is formulated as a Mixed-Integer Linear Programming (MILP) model. It considers the following variables.

- $x_{ij}$ ,  $i \in A^{WS}$ ,  $j \in E$  : Equal to 1 if ambulance  $i$  is dispatched to emergency  $j$  at  $t$  and 0 otherwise.
- $y_{il}$ ,  $i \in A$ ,  $l \in B$  : Equal to 1 if ambulance  $i$  is assigned to base  $l$  at  $t + 1$  and 0 otherwise.
- $r_{ij} \geq 0$ ,  $i \in A^{WS}$ ,  $j \in E$  : Extra response time for ambulance  $i$  to reach emergency  $j$ .
- $u_j \in \{0, 1\}$ ,  $j \in E$  : Equal to 1 if emergency  $j$  is uncovered and 0 otherwise.

- $z_i \in \{0, 1\}$ ,  $i \in A^B \cup A^R$ : Equal to 1 if the target base of ambulance  $i$  at  $t + 1$  is not  $b_i$  and 0 otherwise.
- $v_{lk} \in \{0, 1\}$ ,  $k \in K$ ,  $l \in B$ : Equal to 1 if base  $l \in B$  has, at least, one ambulance of type  $k \in K$  assigned at  $t + 1$  (currently on the road with  $l$  as target base or located at base  $l$ ) and 0 otherwise.
- $w_{qk} \in \{0, 1\}$ ,  $k \in K$ ,  $q \in Z$ : Equal to 1 if no ambulance of type  $k$  can be dispatched to sub-zone  $q$  at  $t + 1$  and 0 otherwise.
- $\delta_{lqk}^{min} \in \{0, 1\}$ ,  $l \in B$ ,  $k \in K$ ,  $q \in Z$ : Equal to 1 if, at  $t + 1$ , the minimum travel time between sub-zone  $q$  and an available ambulance of type  $k$  is obtained with an available ambulance from base  $l$  and 0 otherwise.

The objective is to minimize: system response time, emergencies response time, lacking ambulances, uncovered emergencies, extra response times and changes in target bases. The proposed objective function is the sum of six parts considering the previously described features and it is presented as follows:

$$\begin{aligned}
 \min \quad & \sum_{q \in Z} \sum_{k \in K} \lambda_{qk} \left( \sum_{l \in B} d_{lq} \delta_{lqk}^{min} + INF^{Travel\ time} w_{qk} \right) \\
 & + \sum_{j \in E} \sum_{i \in A^{WS}} (\sigma_j + \delta + d_{ij}) x_{ij} \\
 & + \sum_{k \in K} \left[ P_k^{Ambulance} \sum_{j \in E} (1 + \sigma_j) \left( \eta_j^k - \sum_{i \in A^{WS}: \kappa_i = k} x_{ij} \right) \right] \\
 & + P^{Uncovered} \sum_{j \in E} (1 + \sigma_j) u_j \\
 & + P^{Response} \sum_{i \in A^{WS}} \sum_{j \in E} r_{ij} \\
 & + P^{Target} \sum_{i \in A^{Target}} z_i
 \end{aligned} \tag{2}$$

The M-ADRP model includes the following constraints (3)–(17):

$$(1 - u_j) \leq \sum_{i \in A^{WS}} x_{ij} \leq \sum_{k \in K} \eta_j^k (1 - u_j), \quad \forall j \in E^{\overline{Partial\ Covered}} \tag{3}$$

$$u_j = 0, \quad \forall j \in E^{Partial\ Covered} \tag{4}$$

$$\sum_{i \in A^{WS}: \kappa_i = k} x_{ij} \leq \eta_j^k, \quad \forall j \in E, k \in K \tag{5}$$

$$z_i = 1 - \left( y_{ib_i} + \sum_{j \in E} x_{ij} \right), \forall i \in A^B \cup A^R \quad (6)$$

$$z_i = 0, \forall i \in (A^B \cup A^R) \setminus \{A^{Target}\} \quad (7)$$

$$\sum_{i \in A^{Target}} z_i \leq Target^{MAX} \quad (8)$$

$$\sum_{j \in E} x_{ij} + \sum_{l \in B} y_{il} = 1, \forall i \in A^{WS} \quad (9)$$

$$y_{ih_i} = 1, \forall i \in A^{\overline{WS}} \quad (10)$$

$$\sigma_j + (\delta + d_{ij}) x_{ij} \leq R_j^{MAX} + r_{ij} + M^{Travel\ time} (1 - x_{ij}), \forall i \in A^{WS}, j \in E \quad (11)$$

$$v_{lk} \leq \sum_{i \in A^{WS}: k \in \kappa_i} y_{il} \leq |A| v_{lk}, \forall l \in B, k \in K \quad (12)$$

$$d_{lq} \delta_{lqk}^{min} \leq d_{l'q} + M^{Travel\ time} (1 - v_{l'k}), \forall l, l' \in B, q \in Z, k \in K \quad (13)$$

$$\delta_{lqk}^{min} \leq v_{lk}, \forall l \in B, q \in Z, k \in K \quad (14)$$

$$\sum_{l \in B} \delta_{lqk}^{min} + w_{qk} = 1, \forall q \in Z, k \in K \quad (15)$$

$$x_{ij}, y_{il}, u_j, v_{lk}, z_{i'}, w_{qk}, \delta_{lqk}^{min} \in \{0, 1\},$$

$$\forall i \in A, i' \in A^B \cup A^R, j \in E, l \in B, q \in Z, k \in K \quad (16)$$

$$r_{ij} \geq 0, \forall i \in A^{WS}, j \in E \quad (17)$$

The model constraints ensure the following.

An emergency is uncovered if no ambulance is dispatched to it in the current time period as defined in Constraints (3). Only those emergencies with no ambulance previously assigned can be uncovered at the current period as it is defined in constraints (4).

The number of dispatched ambulances per type is, at most, the required number of ambulances of that type for each emergency (Constraints (5)).

To define which ambulances change their target bases, we consider constraints (6) and (7). The maximum number of allowed changes in the target bases at the same period is considered in Constraint (8) using the input parameter  $Target^{MAX}$ .

Ambulances' status is considered in Constraints (9) which ensure that each available ambulance within the working shift is dispatched to an emergency or assigned to a base; and in Constraints (10) which ensure that ambulances that finished their working shift but are still available must return to their home base, where their status is set to unavailable.

Constraints (11) connect the response time of an emergency (which includes the accumulated waiting time) with the extra response times.

Concerning system's response times, Constraints (13), (14) and (15) determine the minimum travel time for each ambulance type to reach each sub-zone in the following period. To do that it is necessary to know if each base has, at least, one ambulance of each type assigned in the following period which is done through constraints (12).

Finally, constraints (16) and (17) define the variables' domains.

The MILP model formulated in this section was solved using CPLEX 12.7.1 Concert Technology for C++.

## 4.2 Heuristic

A two-phase optimization method is considered. Phase 1 and 2 considers dispatching and relocation decisions, respectively. Phase 2 is solved considering decisions taken in phase 1. Preparedness is embedded in both phases since the main goal is to ensure a good system's coverage. The solution's cost is calculated according to the objective function (2) taking into account its six parts, each one referring to a different feature as previously described when the objective function was defined.

Algorithm 1 describes the heuristic phase 1. The dispatching decision is taken considering the formula presented in line 6. It considers both preparedness and travel time. This means that an available ambulance that takes slightly more time than the closest one can be dispatched if it ensures an improvement in the system's capability to handle future emergencies.

For the phase 2, we consider a *pilot-method* (see [38]) with the pilot heuristic described in Algorithm 2 as a sub-routine.

Both parts of the heuristic proposed in this section (Algorithms 1 and 2) were implemented in C++.

---

### Algorithm 1: Heuristic phase 1—dispatching decisions at $t$

---

```

1  $DISP = \emptyset$  and  $Cost^{DISP} = 0$ 
2 for emergencies  $j \in E$  do
3   for ambulance types  $k \in K$  do
4     for ambulances  $i$  of type  $k$  needed,  $i \in \{1, \dots, \eta_j^k\}$  do
5       Calculate  $p_{(A^{WS} \setminus \{i\})K}^{(t+1)}$ 
6       Dispatch available ambulance  $i'$  such that  $i' = \operatorname{argmax}_{i \in [A^{WS} \setminus DISP : \kappa_i = k]} \frac{p_{(A^{WS} \setminus \{i\})K}^{(t+1)}}{1 + d_{ij}}$ 
7        $DISP = DISP \cup \{i'\}$ 
8       Update  $COST^{DISP}$  (parts 2 and 5 of objective function (2))
9     Update  $COST^{DISP}$  (part 3 of objective function (2))
10  Update  $COST^{DISP}$  if  $j$  is uncovered (part 4 of objective function (2))
```

---

---

**Algorithm 2:** Pilot heuristic—relocation decisions at  $t$ 


---

```

1 Initial ambulance:  $i$ 
2 while [ there are available ambulances to analyze and
3 (number of ambulances that changed target base  $\leq Target^{MAX}$  or there are ambulances at
   an emergency or at a hospital to return to a base) ] do
4   I) Choose base  $l$  to assign ambulance  $i$ 
5   if ambulance  $i$  is at an emergency or at a hospital then
6     Choose between all bases  $l \in B$ , the one that maximizes  $p_{A^{WS}K}^{(t+1)}$ 
7   else
8     if  $i \in A^{Target}$  then
9       Choose between maintaining the current assigned base or changing to the closest
        base, the one that maximizes  $p_{A^{WS}K}^{(t+1)}$ 
10    else
11      Keep the current target base for ambulance  $i$ 
12  II) Choose the following ambulance  $i^*$  to analyze
13   $i^*$  is the furthest available ambulance from  $i$ 
14   $i = i^*$ 
15  III) Update solution cost and other info
16   $Cost_t$  is calculated as the sum of the parts 1 and 6 of the objective function (2)

```

---

### 4.3 Approaches Embedded in a Decision-Support System

To help EMS managers in the decision-making process, we use Geographic Information Systems (GIS) to develop a decision-support tool which embeds the proposed approaches in Sects. 4.1 and 4.2. It was developed using a combination of customization and plugins available through QGIS (a GIS software). A map visualization that analyzes ambulances' movements on the map and the emergencies' location was incorporated. This tool was designed for the Portuguese EMS case, but it can be used for similar EMS systems if the corresponding data is loaded.

## 5 Tests and Analysis

We simulate **50 instances** throughout the day with the following features.

- Emergencies are generated through a *Non-Homogeneous Poisson Process*, using historical data about the average number of occurrences in Lisbon (Fig. 5).
- Emergencies location is random throughout the Lisbon area.
- Emergencies of type I occur with a probability of 0.12, and the real EMS resources are considered.
- In 20% of the emergencies, the patient is not transported to a hospital.
- At the emergency site, each ambulance spends between 1 and 10 min.

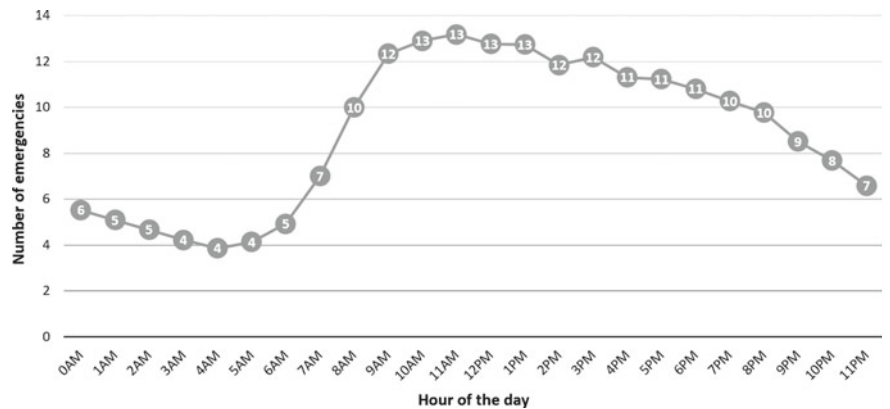


Fig. 5 Average number of emergencies per hour in Lisbon (2nd half of 2016)

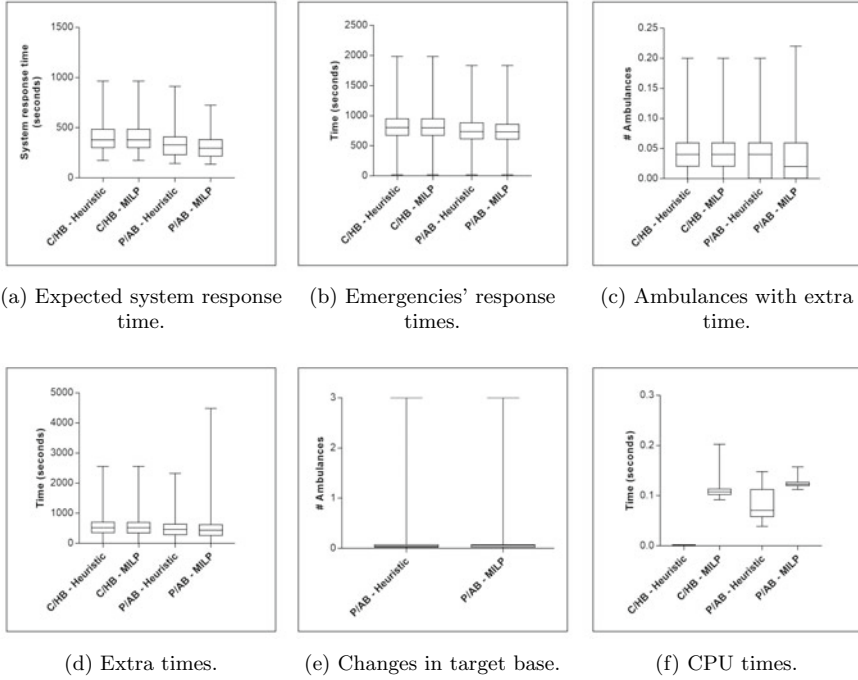
Table 4 Objective function penalizations

|                   |                |
|-------------------|----------------|
| $p^{Target}$      | 3              |
| $p^{Response}$    | 10             |
| $p^{Uncovered}$   | 864000         |
| $p_k^{Ambulance}$ | {86400, 43200} |

- An available ambulance can only change to a different target base after 1 hour from the last change.
- At most, 3 target base changes can occur at the same time.
- Real distances on Lisbon’s roads are considered given by a GIS.
- We consider the objective function penalizations presented in Table 4. The value chosen for  $p^{Target}$  is highly related to the first part of the objective function since changing ambulances’ target base is only worth it when the system’s response time improves for the following time period. Knowing that the average number of emergency calls for all sub-zones throughout the day is approximately 0.03 per min, an improvement of  $3/0.03 = 100$  s in the travel time between an available ambulance and a sub-zone justifies a change in an ambulance target base.  $p^{Response}$  has a value 10 per extra second needed. A penalization of  $86400 \times 10 = 864000$  (being 86400 the seconds of a day) is considered to  $p^{Uncovered}$  which has the highest value since all emergencies must be served. Concerning lacking ambulances, a penalization of 86400 is set to each lacking ambulance of type I and a penalization of 43200 is set to each lacking ambulance of type II. This means that having 10 lacking ambulance of type I or having 20 lacking ambulances of type II is equivalent to an uncovered emergency.

Tests presented in this section were performed on an Intel R Core™ i7-5500U CPU @ 2.40GHz, 8GB RAM.

We compare the proposed strategy (P/AB) with the current Portuguese EMS strategy (C/HB) defined as follows.



**Fig. 6** Simulation main KPIs for M-ADRP

### • Proposed strategy (P/AB)

- *Dispatching*: Policy P (preparedness).
- *Relocation*: Policy AB (any base).

### • Portuguese EMS strategy (C/HB)

- *Dispatching*: Policy C (closest base).
- *Relocation*: Policy HB (home base).

Figure 6 shows the average main key-performance indicators (KPIs) for each minute of simulation considering the experiments previously described.

Figure 6a presents the expected system response time and Fig. 6b presents the average values obtained for response times to reach the emergencies. Figure 6a shows that the P/AB strategy is better than the C/HB strategy. Furthermore, the P/AB-MILP strategy is the one that obtains the best value for the expected system response time. The minimization of the expected system response time is reflected in the emergencies' response times which are lower considering the P/AB strategy.

**Table 5** Comparing C/HB-MILP and P/AB-MILP strategies

|                                                  |               |       |
|--------------------------------------------------|---------------|-------|
| Emergencies' response time ( <i>in seconds</i> ) | Difference    | 80.18 |
|                                                  | % improvement | 10%   |
| Ambulances with extra time                       | Difference    | 0.007 |
|                                                  | % improvement | 16%   |
| Extra time per ambulance ( <i>in seconds</i> )   | Difference    | 70.26 |
|                                                  | % improvement | 13%   |

Figure 6c and d show that, although providing a higher interval for ambulances with extra time and extra times KPIs, the P/AB-MILP strategy is the one with an average lower value for these two KPIs.

Changes in target base can only be compared between P/AB-Heuristic and P/AB-MILP strategies and are analyzed in Fig. 6e. There is no evidence to prefer one of the approaches to minimize this KPI.

Concerning computational times, Fig. 6f shows that C/HB-Heuristic can obtain a solution almost immediately. Although P/AB-MILP takes more time to obtain a solution, it is still obtained in less than one second.

To analyze the impact in the emergencies' response of the proposed strategy, Table 5 presents the average improvement of P/AB over the current Portuguese EMS strategy C/HB using MILP. We obtain an improvement of 80.18 s in the emergencies' response time, which means that ambulances can reach emergencies about one and half minutes earlier than it currently happens. The average number of ambulances with extra time improved 16% and the extra time per ambulance reduces in more than 1 min. These results highlight the proposed approach's potential as it improves the obtained values through the current Portuguese EMS strategy for the most relevant KPIs.

## 6 Conclusion

Dispatching and relocation decisions are considered for the multiple ambulance type dispatching and relocation problem. Two different optimization approaches are studied: a MILP model and a two-phase heuristic. To help the EMS managers in the decision-support process of dispatching and relocation tasks, we develop a GIS-based tool, which embeds the proposed approaches.

The proposed strategy uses a preparedness rule for dispatching and allows relocations to any base—P/AB. P/AB strategy is compared with the current Portuguese strategy, which dispatches the closest available ambulance and relocates ambulances to their home bases—C/HB. The potential of the P/AB strategy under the MILP approach is highlighted since it obtains, in general, better results for the main KPIs such as the response times.

**Acknowledgements** This research is partially supported by FCT, under project UIDB/04561/2020.

## References

1. Aboueljinane, L., Sahin, E., Jemai, Z.: A review on simulation models applied to emergency medical service operations. *Comput. Ind. Eng.* **66**(4), 734–750 (2013)
2. Alanis, R.: Emergency medical services performance under dynamic ambulance redeployment. PhD thesis, University of Alberta (2012)
3. Alanis, R., Ingolfsson, A., Kolfal, B.: A markov chain model for an EMS system with repositioning. *Prod. Oper. Manag.* **22**(1), 216–231 (2013)
4. Andersson, T., Värbrand, P.: Decision support tools for ambulance dispatch and relocation. *J. Operat. Res. Soc.* **58**(2), 195–201 (2007)
5. Aringhieri, R., Bruni, M.E., Khodaparasti, S., van Essen, J.T.: Emergency medical services and beyond: Addressing new challenges through a wide literature review. *Comput. Oper. Res.* **78**, 349–368 (2017)
6. Bélanger, V., Lanzarone, E., Ruiz, A., Soriano, P.: The ambulance relocation and dispatching problem. Technical Report, Interuniversity Research Center on Enterprise Networks, Logistics and Transportation, CIRRELT-2015-59 p 17 (2015)
7. Bélanger, V., Ruiz, A., Soriano, P.: Recent optimization models and trends in location, relocation, and dispatching of emergency medical vehicles. *Eur. J. Oper. Res.* **272**(1), 1–23 (2019)
8. Bélanger, V., Lanzarone, E., Nicoletta, V., Ruiz, A., Soriano, P.: A recursive simulation-optimization framework for the ambulance location and dispatching problem. *Eur. J. Oper. Res.* **286**(2), 713–725 (2020)
9. Billhardt, H., Lujak, M., Sánchez-Brunete, V., Fernández, A., Ossowski, S.: Dynamic coordination of ambulances for emergency medical assistance services. *Knowl.-Based Syst.* **70**, 268–280 (2014)
10. Brotcorne, L., Laporte, G., Semet, F.: Ambulance location and relocation models. *Eur. J. Oper. Res.* **147**(3), 451–463 (2003)
11. Carvalho, A.S., Captivo, M.E., Marques, I.: Integrating the ambulance dispatching and relocation problems to maximize system's preparedness. *Eur. J. Oper. Res.* **283**(3), 1064–1080 (2020)
12. Church, R., ReVelle, C.: The maximal covering location problem. *Pap. Reg. Sci.* **32**(1), 101–118 (1974)
13. Comber, A.J., Sasaki, S., Suzuki, H., Brunson, C.: A modified grouping genetic algorithm to select ambulance site locations. *Int. J. Geogr. Inf. Sci.* **25**(5), 807–823 (2011)
14. Goldberg, J.B.: Operations research models for the deployment of emergency services vehicles. *EMS Manag. J.* **1**(1), 20–39 (2004)
15. Ibri, S., Nourelfath, M., Drias, H.: A multi-agent approach for integrated emergency vehicle dispatching and covering problem. *Eng. Appl. Artif. Intell.* **25**(3), 554–565 (2012)
16. Jagtenberg, C.J., Bhulai, S., van der Mei, R.D.: An efficient heuristic for real-time ambulance redeployment. *Operat. Res. Health Care* **4**, 27–35 (2015)
17. Jagtenberg, C.J., Bhulai, S., van der Mei, R.D.: Dynamic ambulance dispatching: is the closest-idle policy always optimal? *Health Care Manag. Sci.* **20**(4), 517–531 (2016)
18. Jagtenberg, C.J., van den Berg, P.L., van der Mei, R.D.: Benchmarking online dispatch algorithms for emergency medical services. *Eur. J. Oper. Res.* **258**(2), 715–725 (2017)
19. Karasakal, E., Silav, A.: A multi-objective genetic algorithm for a bi-objective facility location problem with partial coverage. *TOP* **24**(1), 206–232 (2015)
20. Kergosien, Y., Bélanger, V., Soriano, P., Gendreau, M., Ruiz, A.: A generic and flexible simulation-based analysis tool for EMS management. *Int. J. Prod. Res.* **53**(24), 7299–7316 (2015)
21. Kolesar, P., Walker, W.E.: An algorithm for the dynamic relocation of fire companies. *Oper. Res.* **22**(2), 249–274 (1974)
22. Lee, S.: The role of preparedness in ambulance dispatching. *J. Operat. Res. Soc.* **62**(10), 1888–1897 (2011)
23. Lee, S.: A new preparedness policy for EMS logistics. *Health Care Manag. Sci.* **20**(1), 105–114 (2017)

24. Li, X., Zhao, Z., Zhu, X., Wyatt, T.: Covering models and optimization techniques for emergency response facility location and planning: a review. *Math. Methods Oper. Res.* **74**(3), 281–310 (2011)
25. Maxwell, M.S., Restrepo, M., Henderson, S.G., Topaloglu, H.: Approximate dynamic programming for ambulance redeployment. *INFORMS J. Comput.* **22**(2), 266–281 (2010)
26. Nasrollahzadeh, A.A., Khademi, A., Mayorga, M.E.: Real-time ambulance dispatching and relocation. *Manuf. Serv. Oper. Manag.* **20**(3), 467–480 (2018)
27. Owen, S.H., Daskin, M.S.: Strategic facility location: a review. *Eur. J. Oper. Res.* **111**(3), 423–447 (1998)
28. Rajagopalan, H.K., Vergara, F.E., Saydam, C., Xiao, J.: Developing effective meta-heuristics for a probabilistic location model via experimental design. *Eur. J. Oper. Res.* **177**(1), 83–101 (2007)
29. Reuter-Oppermann, M., van den Berg, P.L., Vile, J.L.: Logistics for emergency medical service systems. *Health Syst.* **6**(3), 187–208 (2017)
30. ReVelle, C.: Review, extension and prediction in emergency service siting models. *Eur. J. Oper. Res.* **40**(1), 58–69 (1989)
31. ReVelle, C., Bigman, D., Schilling, D., Cohon, J., Church, R.: Facility location: a review of context-free and EMS models. *Health Serv. Res.* **12**(2), 129–146 (1977)
32. Sasaki, S., Comber, A.J., Suzuki, H., Brunsdon, C.: Using genetic algorithms to optimise current and future health planning - the example of ambulance locations. *Int. J. Health Geogr.* **9**(1), 4 (2010)
33. Schmid, V.: Solving the dynamic ambulance relocation and dispatching problem using approximate dynamic programming. *Eur. J. Oper. Res.* **219**(3), 611–621 (2012)
34. Toregas, C., Swain, R., ReVelle, C., Bergman, L.: The location of emergency service facilities. *Oper. Res.* **19**(6), 1363–1373 (1971)
35. Van Barneveld, T.C., Bhulai, S., van der Mei, R.D.: The effect of ambulance relocations on the performance of ambulance service providers. *Eur. J. Oper. Res.* **252**(1), 257–269 (2016)
36. Van Barneveld, T.C., van der Mei, R.D., Bhulai, S.: Compliance tables for an EMS system with two types of medical response units. *Comput. Oper. Res.* **80**, 68–81 (2017)
37. Van Buuren, M., Jagtenberg, C., van Barneveld, T., van Der Mei, R., Bhulai, S.: Ambulance dispatch center pilots proactive relocation policies to enhance effectiveness. *Interfaces* **48**(3), 235–246 (2018)
38. Voß, S., Fink, A., Duin, C.: Looking ahead with the pilot method. *Ann. Oper. Res.* **136**(1), 285–302 (2005)
39. Zhen, L., Wang, K., Hu, H., Chang, D.: A simulation optimization framework for ambulance deployment and relocation problems. *Comput. Ind. Eng.* **72**, 12–23 (2014)

# Automated Radiotherapy Treatment Planning Optimization: A Comparison Between Robust Optimization and Adaptive Planning



Hugo Ponte, Humberto Rocha, and Joana Dias

**Abstract** Radiotherapy treatments must guarantee a proper coverage of the volume to treat (Planning Target Volume—PTV) while sparing all the organs at risk (OAR). As the treatment is usually fractionated during a given planning horizon, replanning is many times needed. However, replanning has some shortcomings, due to the need of acquiring new medical images, and additional time needed for creating new treatment plans. This work studies how the use of robust planning compares with replanning, considering the impact on PTV coverage and OAR sparing. Four different treatment planning approaches are compared. The conventional approach keeps the same treatment plan during the whole treatment time. A new robust approach is tested, where possible scenarios for the PTV evolution are also considered when creating the initial treatment plan. A third approach considers replanning once halfway of the treatment time. A fourth approach mimics what is usually known as the “plan of the day”. These approaches were applied to a head-and-neck cancer case and compared by Monte Carlo simulation. The robust approach originated better treatment plans than the conventional approach and it can be a competitive alternative to replanning.

**Keywords** Radiotherapy · Treatment planning optimization · Adaptive planning · Inverse optimization

---

H. Ponte

Faculty of Sciences and Technology, University of Coimbra, Coimbra, Portugal

H. Rocha

Faculty of Economics and CeBER, University of Coimbra, Coimbra, Portugal

e-mail: [hrocha@fe.uc.pt](mailto:hrocha@fe.uc.pt)

J. Dias (✉)

Faculty of Economics and INESC Coimbra, University of Coimbra, Coimbra, Portugal

e-mail: [joana@fe.uc.pt](mailto:joana@fe.uc.pt)

## 1 Introduction

According to the global cancer observatory, the number of new cases of cancer worldwide in 2020 was almost 19.3 million, 4.4 million of those being in Europe [1]. Many treatments can be provided according to the type and stage of the cancer, including radiotherapy, chemotherapy, surgery, stem cell transplant, immunotherapy, among others. Radiation therapy (or radiotherapy) delivers high doses of radiation aiming to kill cancer cells by damaging their DNA.

External radiotherapy is delivered with the patient laying on a couch that can rotate. Radiation is generated by a linear accelerator, or linac, placed on a gantry that can rotate around a central axis parallel to the couch. When the couch is fixed at  $0^\circ$  the beam directions are called coplanar. Rotation of the gantry and the couch allows noncoplanar beams irradiating the tumour from almost any direction around the patient, except for directions that would cause collisions between the gantry and the couch/patient [2]. The main objective of radiotherapy treatment planning is to be able to calculate a treatment plan that irradiates adequately the volumes to treat while, at the same time, sparing as most as possible all the surrounding organs, keeping their functionality. The main focus of this work will be Intensity Modulated Radiation Therapy (IMRT), where it is possible to modulate the radiation delivered so that it better conforms with the volumes that should be treated.

The treatment is usually fractionated, being delivered daily during several weeks. This fractionation aims at malignant cell destruction while preserving healthy tissue. The fractionation means that the volumes to be treated, as well as the volumes that correspond to organs that should be spared, are not exactly the same during the whole treatment time. Changes in these volumes could be accommodated if new medical images were acquired and new treatment plans were calculated, but this alternative comes with a cost, as discussed later on. The impact of changes in the volumes to treat and different approaches to take these changes into account are the main focus of this work. Four different approaches will be compared:

- Only one treatment is planned, and no other replanning is done during the course of treatment. This initial treatment plan is calculated considering only the existing structures of interest as defined in the planning medical image (computed tomography scan – CT).
- Only one treatment plan is considered during the course of treatment, but this plan is calculated considering not only the original structures of interest but also additional structures that are created and that represent possible changes to the PTV (it is, in some sense, a robust plan taking into consideration different potential scenarios for the PTV).
- Replanning is performed once during the course of treatment.
- Replanning is performed more often than the previous approach, mimicking the concept of the “plan of the day”.

## 2 Intensity Modulated Radiotherapy Treatment Planning

IMRT uses a multi-leaf collimator (MLC) to modulate the beams [3]. The MLC consists of pairs of leaves that can move side by side creating a variety of field openings. The movement of the leaves originates the segmentation of each beam into a set of smaller imaginary beams with independent fluence (intensities), called beamlets, delivering a non-uniform radiation field to the patient [2, 4].

The planning of a given treatment begins by considering a treatment CT scan and the medical prescription. In the CT scan all the structures, OARs and PTVs must be well mapped. The OARs are the organs that may be exposed to radiation, while the PTV represents the volume of the known tumour including microscopic spread plus a margin around this volume. Margins around target structures exist to compensate for eventual inaccuracies. All the structures of interest are discretized into voxels. Voxels are volumetric pixels where the length and width of each voxel depends on the resolution and spacing between the CT images [5].

Absorbed dose is the energy deposited per unit of mass of tissue and it is expressed in Gray (Gy). The medical prescription, defined by the medical oncologist, is patient dependent and defines a set of dosimetric goals and constraints, for instance, maximum dose deposited in OARs and minimum dose deposited in PTVs. The definition of OAR constraints aims at preserving the organ functionality. Usually, two different sets of OARs are defined: serial and parallel organs. Serial organs are the ones that, if even only a small part of the OAR volume is damaged, the OAR functionality is jeopardized. In this case, usually the constraints define a maximum dose that cannot be surpassed at any voxel. Parallel organs are the ones that keep their functionality even if only a small part is damaged. In this case, the constraints usually define a mean dose threshold, taking into consideration the dose deposited in all the OAR's voxels. It is possible to define dose-volume constraints stating that only a certain fraction of a structure can be exposed to a dose value higher (lower) than the upper (lower) threshold.

Treatment planning is typically organized in 3 stages: (1) Beam angle optimization problem; (2) Intensity/Fluence map optimization problem; (3) Realization Problem. In the first stage, decisions that define the positioning of the patient and the irradiation directions are made. The third stage considers the optimization of the movement of the MLC leaves so that the fluence map calculated in step 2 is delivered. In this work we will focus on the second stage, the optimization of the fluence map.

### 2.1 Fluence Map Optimization

The fluence map optimization (FMO) problem consists of finding the optimal beamlet intensities (weights) for the previously chosen beam angles. The dose for each voxel is calculated using the superposition principle, where the contribution of each beamlet is considered. A dose matrix,  $D$ , is obtained from the unitary fluence of each beamlet.

It has information on the absorbed dose in each voxel, from each beamlet, considering radiation intensity equal to 1. In this matrix, each row corresponds to one voxel, and each column corresponds to one beamlet. The dose received by voxel  $i$ ,  $d_i$ , can be represented as the sum of the contribution of each beamlet,  $d_i = \sum_{j=1}^{N_b} D_{ij} w_j$ , where  $N_b$  is the total number of beamlets,  $w_j$  is the intensity of the beamlet  $j$  and  $D_{ij}$  represents the elements in the  $i$ th row and  $j$ th column of matrix  $D$  (absorbed dose in voxel  $i$  from beamlet  $j$ , considering beamlet  $j$  has unitary radiation intensity). The matrix  $D$  is a very large matrix, since the total number of voxels,  $N_v$ , can reach tens of thousands, which makes the fluence optimization hard.

In this work, FMO resorts to the minimization of a quadratic objective function with no additional constraints. For each and every voxel, the dose received is compared with a given threshold. The objective function will penalize any difference between the received dose and the desired/allowed dose for each voxel. These differences are squared in order to obtain a convex problem that is easy and fast to solve (see Eq. 1).

$$\min_{w \geq 0} \sum_{i=1}^{N_v} \left[ \underline{\lambda}_i \left( T_i - \sum_{j=1}^{N_b} D_{ij} w_j \right)^2_{+} + \overline{\lambda}_i \left( \sum_{j=1}^{N_b} D_{ij} w_j - T_i \right)^2_{+} \right] \quad (1)$$

where  $T_i$  is the desired/allowed dose for voxel  $i$ ,  $\underline{\lambda}_i$  and  $\overline{\lambda}_i$  are the penalties of underdose and overdose of the voxel  $i$ , and  $(\cdot)_{+} = \max\{0, \cdot\}$  [5].

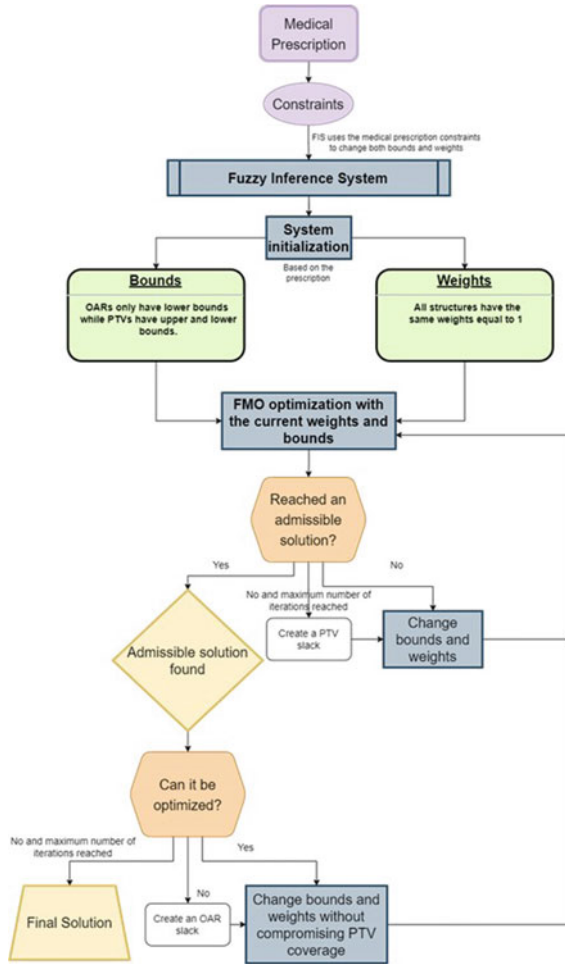
It is important to note that this objective function has no clinical meaning, even though it is related with the dose volume constraints defined by the medical prescription. The parameters used in the objective function, namely weights and bounds, can be freely changed, being possible to steer the search so that the optimal solution for this problem is indeed an admissible solution from a clinical point-of-view (complying with all the restrictions defined in the medical prescription). This is actually what is done by the treatment planner. The planner is assisted by a dedicated software, called Treatment Planning System (TPS), and is asked to define several parameters associated with the structures of interest, like lower and upper dosimetric bounds or weights. If the dose distribution achieved with the current parameters meets the medical prescription the procedure finishes. Otherwise, one or more parameters are updated, and a new treatment plan is calculated. In this lengthy trial-and-error procedure, the planners typically resort to their own experience, and the resulting plans may not even be the best ones in terms of sparing critical organs or the achievement of proper PTV irradiation [5]. As a planner may need hours or even days to reach a high-quality treatment plan, many efforts have been made trying to make treatment planning partially or fully automated. Zarepisheh et al. propose a treatment planning optimization based on the DVH curves of a reference plan, where voxel weights used in the FMO objective function are automatically updated by projecting the current dose distribution on the Pareto surface of the problem, considering the corresponding gradient information [6]. Jia et al. propose a treatment plan procedure in an OAR 3D dose distribution prediction [7]. Instead of having a medical prescription guiding the

FMO, these predictions are the dosimetric references for OAR objectives. The FMO objective function is not dynamically updated taken into consideration the current solution dosimetric achievements. Schipaanboord et al. propose a fully automated treatment planning procedure for robotic radiotherapy [8] and Bijman et al. for MRI-Linac (uses magnetic resonance imaging to monitor the target area at the same time the treatment is being delivered) applied to rectal cancer [8, 9]. This method requires that a fixed wish-list is previously defined to be used for all the patients with the same tumour location. The need to fix a priori all the parameters in the wish-list can be a disadvantage for some patients that have specific situations that deviate from the most common cases. A new fully automated FMO procedure, where all the objective function parameters (weights and bounds) are tuned based on fuzzy inference systems (FIS) was presented in Dias et al. [10]. The authors propose the use of fuzzy inference systems (FIS) to mimic the planner's procedure in an automated, efficient and optimized way by iteratively changing both weights and bounds in Eq. (1) based on a reasoning similar to the human planners. Lower and upper bounds are initialized based on the medical prescription. Upper and lower bounds are assigned to PTVs while OARs only need upper bounds. OARs only have dose volume constraints that consider upper bounds, whilst PTVs have usually dose volume constraints that consider both upper and lower bounds (it is necessary to guarantee, for instance, that 95% of the volume receives at least 95% of the prescribed dose but, at the same time, to guarantee that no voxel receives more than 107% of the prescribed dose).

The planners' reasoning can be expressed by a set of simple rules in natural language: if, for a certain structure, the deviation between the received dose and the prescribed dose is low/medium/large, then the respective upper/lower bound should be slightly/medium/greatly increased/decreased. The further away the structure of interest is from what is desired, the more pronounced the change of the respective parameter should be. Fuzzy numbers allow for the mathematical representation of these concepts, like low, medium or large. Weights are changed in a similar way as bounds. Figure 1 depicts the FIS flowchart. In this work all treatment plans are calculated by resorting to this fuzzy approach. Calculating fluence maps by using the same algorithm makes the comparisons among the different approaches fair. As this is a truly automated approach, it also makes it feasible to calculate fluence maps even when an increased number of structures is considered, as is the case of the robust optimization approach described.

### 3 Uncertainties and Treatment Planning

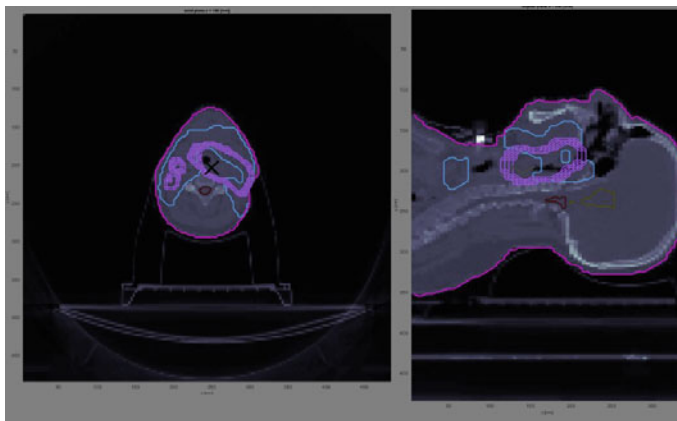
The planning CT represents the patient on the day the medical image was acquired. However, there is a time lag between this CT being acquired and the treatment being initiated and, as the treatment is divided into a set of daily fractions, the patient's anatomy changes as the treatment progresses, possibly leading to discrepancies between planned and delivered dose distributions. These discrepancies can lead to a decrease in the probability of tumour control, or increased probability of



**Fig. 1** FIS flowchart

OAR complications. Adaptive radiotherapy (ART) was introduced as a method to reduce the effects of patient related uncertainties: it considers the replanning of the treatment, with acquisition of new medical images, taking into account the tumour evolution and other possible anatomical changes [11]. The increase of the number of times a replanning is performed converges to the concept of the “plan of the day”. In this case, based on a daily CT, the treatment can be replanned also daily and immediately before the treatment is delivered. The technique allows for better PTV coverage and organ sparing [12]. However, replanning comes at a cost: increased total radiation delivered to the patient due to the increased number of CTs that are acquired. Furthermore, whenever a CT is acquired, all the volumes of interest must be delineated (which is not yet a totally automated procedure), and a new plan must

be calculated, which takes time. This can have an important impact in the treatment delivery workflow, and also on the total number of patients that can be treated daily. Instead of calculating an initial treatment plan that considers only the information given by the planning CT, we resort to the concept of robust optimization and consider additional structures that represent different scenarios for the positioning and volume of the PTV. Although all the structures of interest are expected to change during treatment, in this work a simplification has been assumed and only PTV was considered as being changed during treatment, keeping unaltered all the other structures. We decided to only consider the case where the target volume decreases in size, which is the majority of cases [13]. In our robust treatment planning approach, 14 additional auxiliary structures are considered in the FMO. These structures are copies of the PTV, but in a different position considering shifts in the  $x$ ,  $y$  and  $z$  axes: one for each quadrant bisector and two for each axis for both positive and negative directions. This means that we end up with 14 more auxiliary structures for the PTV: right and left shifts for each axe alone (6 auxiliary structures), right and left shifts considering simultaneous two axes:  $(x, y)$ ,  $(x, z)$ ,  $(y, z)$ , and right and left shifts considering the three axes at the same time. All these 14 structures have exactly the same size as the original PTV, as shown in Fig. 2. In this figure it is possible to see the delineation of these additional structures (in purple), from two different visual plans. The auxiliary structures were considered in the initial treatment plan and, even though they are not real, they are treated equally as the PTV. The goal of the inclusion of these structures in the optimization loop is to guarantee a proper coverage of the PTV considering that it can be in different positions and with a different volume than its original presentation in the planning CT. Although the scenarios consider a shift of the volume only, and not its decrease, they make the fluence map optimization



**Fig. 2** Images of the head-and-neck cancer case with 14 auxiliary structures around the PTV70. The purple lines shown are defining different volumes around the PTV that correspond to shifts, along the three axes, of the PTV original volume

produce a more homogeneous dose distribution that can be an advantage when the PTV volume changes.

## 4 Computational Experiments

### 4.1 Head-and-Neck Cancer Case

To compare the different treatment planning approaches, a head-and-neck cancer case was used. This case is available in matRad, the open-source treatment planning system for academic research that was used in this work. Treatment planning for head-and-neck cancer cases is usually a very complex procedure. In the case considered, four OARs were included in the treatment plan optimization (right and left parotids, brainstem, spinal cord), in addition to two PTVs with different medical prescriptions that are overlapping. All the voxels that do not belong to any of these structures of interest are assigned to an OAR usually named as Body or Skin. This OAR is also important in treatment planning because it has to be guaranteed that no hotspots appear in the patient outside the structures of interest delineated. OARs like the spinal cord and brainstem are serial organs. Therefore, a maximum dose is established for these organs. The parotids are parallel organs. The tolerance defined for these organs depends on the size of the volume irradiated, therefore mean dose is usually the objective type considered for these organs.

For the PTVs, the usual medical prescription determines that a given percentage of the volume has to receive at least a given dose. In this case, the dose of 95% of the volume ( $D_{95}$ ) is considered, and this value has to be at least 95% of the prescribed dose. If the prescription dose is equal to 70Gy, for instance, then,  $D_{95}$  should be greater than or equal to 66.5Gy, meaning that at least 95% of the PTV volume has to receive at least 66.5 Gy. In this case, one PTV has a medical prescription of 70Gy (named PTV70), the other one of 63 Gy (named PTV63). PTV70 is enclosed in PTV63, being the latter used to guarantee a proper dose gradient in the PTV70 adjacent areas. For this reason, only changes in the PTV70 are being considered. The medical prescription for each of the 7 original structures is described in Table 1.

The beam angle configuration considered was the one usually chosen in clinical practice: seven fixed, equidistant coplanar angles.

### 4.2 Treatment Planning Strategies

Our first approach consisted in obtaining an admissible plan based only on the planning CT. This treatment planning is not accounting for any changes that might happen to the PTV70. The second approach, referred to as *New approach*, also consisted in obtaining an admissible plan based only on the planning CT, but considering an

**Table 1** Medical prescription for a head-and-neck cancer patient in Gray (Gy)

| Structure     | Type of constraint |              | Limit |
|---------------|--------------------|--------------|-------|
| BrainStem     | Maximum dose       | Lower than   | 54    |
| Left parotid  | Mean dose          | Lower than   | 26    |
| Right parotid | Mean dose          | Lower than   | 26    |
| Skin          | Maximum dose       | Lower than   | 80    |
| Spinal cord   | Maximum dose       | Lower than   | 45    |
| PTV63         | $D_{95}$           | Greater than | 59.85 |
| PTV70         | $D_{95}$           | Greater than | 66.5  |
| PTV70         | Maximum dose       | Lower than   | 74.9  |

additional set of auxiliary structures similar to the PTV70 as explained before. This will change the way in which the fluences are optimized and calculated, since it will be necessary to comply with the medical prescription in more than one volume. To acknowledge the possible improvements of adding replanning even in the presence of a robust treatment plan, the previously mentioned method was repeated but with the added feature of doing a replanning halfway through the treatment. This plan is referred to as *Replan*. Lastly, replanning is considered as often as a change in the PTV is simulated, mimicking the *Plan of the day*.

All four approaches used the FIS algorithm to change in an automated way bounds and weights of both PTV and OARs in the FMO procedure.

### 4.3 Plan Quality Assessment

To assess the performance of the four different approaches, we use Monte Carlo (MC) simulation, taking explicitly fractionation into account. During each iteration (that corresponds to one total treatment), the PTV70 will randomly change at pre-determined fractions. All the dosimetric indicators of interest are then calculated for each iteration of the simulation: The dose received by at least 95% of the PTV70 ( $D_{95}$ ), mean dose ( $D_{mean}$ ) for parotids, and maximum dose ( $D_{max}$ ) for spinal cord, brainstem and Skin. These dosimetric indicators are calculated considering the dose deposited in each one of the voxels of the respective structures during each one of the treatment fractions.

### 4.4 Computational Tests

Our tests were run on an Intel Core i7-8750H 2.2GHz. matRad [14], an open-source software that has been developed within the MATLAB environment, was used for

dose calculation. It mimics the use of a commercial TPS, but for academic use only. MATLAB version 2018b was used.

A beam angle configuration of 7 equidistant coplanar angles [0 52 103 154 205 256 307] was considered, meaning that the couch angle will always be 0. This is the most usual configuration used in the clinical practice for similar cases. For treatment quality assessment, the MC simulations considered 50 iterations, each with 35 fractions. A random change in the PTV is considered every 5 fractions. To avoid having any bias in the results due to the random changes generated, all the different approaches were tested considering the same sequence of random numbers, so that exactly the same PTV70 changes were considered in all the comparisons made. This structure will be changed 6 times during the course of each iteration of the simulation. The tumour volume is thus randomly and iteratively changed to simulate asymmetrical size decreases during the treatment course. These changes were calculated as follows: based on the current structure, small and randomly generated displacements in the  $x$ ,  $y$  and  $z$  axes relative to its current position are considered, and auxiliary structures were created. The intersection between the current PTV and one of the new auxiliary structures gives rise to a new PTV that corresponds to an asymmetric reduction of the original structure. Figure 3 illustrates this procedure.

## 4.5 Results

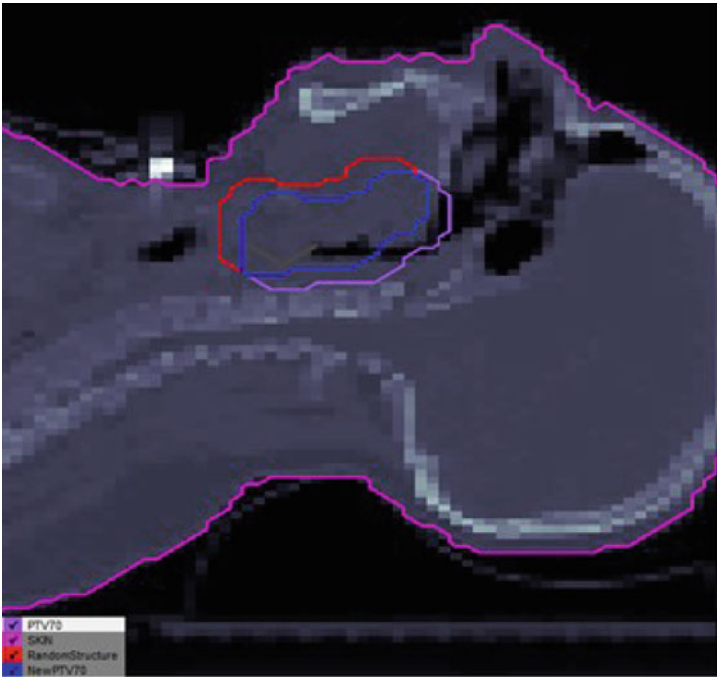
The main objective of this study is to test and compare different treatment planning approaches, introducing a new approach based on auxiliary structures to be used in the initial treatment planning procedure.

It was possible to obtain clinically admissible treatment plans considering all 50 iterations of the MC simulation for the four different approaches tested. It was thus possible to respect both upper and lower bounds defined in Table 1, for all the structures.

The dosimetric values of all OARs for the *Conventional* and *New approach* stay the same throughout the treatment since there is not replan (Table 2). For each and every case and for all OARs it was possible to have values below the maximum threshold defined. For every case all plans respected the PTV70 thresholds established in the medical prescription.

The plans were evaluated regarding the obtained dosimetric values for each structure. For the PTV70,  $D_{95}$  was the considered criteria, for the brainstem and spinal cord  $D_{max}$  was considered and for the parotids  $D_{mean}$  was considered. Table 2 displays the results obtained for the different approaches.

The *Conventional* approach obtained the worst result among the four plans regarding PTV70 coverage ( $D_{95} = 67.45 \pm 0.09$  Gy). On the other hand, it obtained the best dosage value regarding the right parotid ( $D_{mean} = 19.48 \pm 0.0$  Gy), side by side with the *Replan* approach. It also obtained one of the lowest values, statistically identical to the *Replan* approach regarding the left parotid with  $D_{mean} = 21.07 \pm 0.0$  Gy.



**Fig. 3** Generating an auxiliary structure by randomly shifting the PTV70. The intersection of the voxels of the original PTV70 and this auxiliary structure results in a new randomly generated reduced PTV70. This can be seen in this image by looking at the intersection between the purple (original PTV) and the red (randomly generated structure) volumes: the resulting volume is represented in blue

**Table 2** Differences among the 4 plans, for the structures of interest. The best results are highlighted in grey for each structure. Units are in Gray(Gy)

|               |            | Conventional       |      | New approach       |      | Replan             |      | Plan of the day    |      | Z          |
|---------------|------------|--------------------|------|--------------------|------|--------------------|------|--------------------|------|------------|
|               |            | Mean               | SD   | Mean               | SD   | Mean               | SD   | Mean               | SD   |            |
| PTV70         | $D_{95}$   | 67.45 <sup>a</sup> | 0.09 | 68.07 <sup>b</sup> | 0.09 | 67.82 <sup>c</sup> | 0.18 | 68.01 <sup>b</sup> | 0.16 | 204.42***  |
| BrainStem     | $D_{max}$  | 44.71 <sup>a</sup> | 0.0  | 44.55 <sup>a</sup> | 0.0  | 45.01 <sup>b</sup> | 0.46 | 45.26 <sup>b</sup> | 0.49 | 43.86***   |
| Spinal Cord   |            | 40.70 <sup>a</sup> | 0.0  | 41.01 <sup>b</sup> | 0.0  | 40.15 <sup>c</sup> | 0.34 | 40.20 <sup>c</sup> | 0.45 | 107.14***  |
| Left Parotid  | $D_{mean}$ | 21.07 <sup>a</sup> | 0.0  | 22.57 <sup>b</sup> | 0.0  | 21.05 <sup>a</sup> | 0.11 | 21.03 <sup>a</sup> | 0.12 | 4356.84*** |
| Right Parotid |            | 19.48 <sup>a</sup> | 0.0  | 20.08 <sup>b</sup> | 0.0  | 19.48 <sup>a</sup> | 0.08 | 19.53 <sup>c</sup> | 0.09 | 1303.01*** |

*Note* Means with different letters are significantly different at the level of  $\alpha < 0.001$  according to the post-hoc test of Tukey HSD. \*\*\*  $p < 0.001$ . This means that no statistical evidence was found to allow rejecting the hypothesis that the mean values with the same letter are equal. As an example, for PTV70 and with means 68.07<sup>b</sup> and 68.01<sup>b</sup>, the hypothesis of their being statistically equal was not rejected

The *New approach* showed the best results regarding PTV70 coverage with  $D_{95} = 68.07 \pm 0.09$  Gy. It also showed the best OAR sparing regarding the brain stem with  $D_{max} = 44.55 \pm 0.0$  Gy.

The *Replan* approach behaved the best for the spinal cord, and the right parotid with mean dosimetric values of  $D_{max} = 40.15 \pm 0.34$  Gy and  $D_{mean} = 19.48 \pm 0.08$  Gy, respectively. In terms of PTV70 coverage it was slightly worst ( $D_{95} = 67.82 \pm 0.18$  Gy) than the two best approaches, *Plan of the day* and *New approach*. The *Plan of the day* showed, along with *New approach*, the best PTV70 coverage results ( $D_{95} = 68.01 \pm 0.16$  Gy). It also obtained the best results regarding the left parotid sparing, with  $D_{mean} = 21.03 \pm 0.12$  Gy, and regarding the spinal cord, with  $D_{max} = 40.20 \pm 0.45$  Gy.

## 5 Discussion and Conclusions

Adaptative approaches have the advantage of using updated patient information, namely considering the delineation of all the structures of interest, during the course of the treatment to improve treatment quality. Replanning is a very important procedure in radiotherapy, usually resulting in better dosimetric results when compared to the use of a single treatment plan calculated before the beginning of the treatment.

In this work, a new approach is proposed where new auxiliary structures were created and used in the treatment planning procedure. These auxiliary structures enable a robust treatment plan to be calculated from the beginning of the treatment.

Four treatment planning approaches were compared, by evaluating how each one of them behaves regarding the target volume coverage while fulfilling the thresholds of the surrounding organs. Monte Carlo simulation was used to compare the different approaches, considering 50 iterations, each with 35 fractions, during which the PTV was iteratively changed. It was possible to conclude that generating new auxiliary structures around the targeted area revealed enhanced PTV70 coverage when compared to *Conventional* plans. The approach also got improved brainstem sparing. The proposed new methodology behaved as expected regarding the target volume coverage.

Overall, the replanning approaches showed better results in sparing the OARs, while improving the PTV70 coverage. On the other hand, the *New approach* showed the best results improving PTV70 coverage, but generally did not show better results sparing the OARs. The *New approach* and the replanning approaches showed improved target coverage compared to the *Conventional* approach which was the main feature in study. The results proved to be consistent with the existing literature and experiments.

Although we believe the results and conclusions reached are interesting and valuable, it is also important to identify the limitations of this study.

It is important to mention that only changes to the PTV70 were considered, whilst all the other structures were considered unaltered during the course of treatment. In a real situation, all the structures suffer changes, some of which can even be correlated. Head-and-neck tumour cases have many OARs nearby the PTV and, during the treatment, when the PTV diminishes, the tissue around it also moves, moving the surrounding OARs near the target area. These changes were not replicated in the Monte Carlo simulations performed. Real head-and-neck cases also have more OARs that are usually taken into account, depending on the PTV and location, like

the larynx, eyes, oral cavity, other salivary glands, etc. We have only considered 4 OARs. Moreover, only one case was considered, as a proof-of-concept for the new approach proposed. It would be important to repeat these computational tests with a representative set of cases to see if these results are generalisable or not.

In summary, using auxiliary structures proved to be a valid approach regarding PTV coverage. Replanning more often also showed improved PTV coverage in all cases compared to the conventional methodology while maintaining OAR sparing specially regarding the left parotid and the spinal cord. The new methodology using auxiliary structures proved to be a competitive alternative to replanning, avoiding the time-consuming replanning, and sparing the patient to further imaging sessions and all the possible damages associated.

As future work it would be interesting to assess the proposed approach performance for other cancer cases and include beam angle optimization.

**Acknowledgements** This work has been supported by the Fundação para a Ciência e a Tecnologia (FCT) under project grants UIDB/05037/2020, UIDB/00308/2020 and UTA-EXPL/FMT/0079/2019, and FEDER under grant POCI-01-0247-FEDER-047222.

## References

1. Globocan, Cancer facts 2020 (2020). <https://gco.iarc.fr/today/data/factsheets/cancers/39-All-cancers-fact-sheet.pdf>
2. Dias, J., Rocha, H., Ferreira, B., Lopes, M.C.: A genetic algorithm with neural network fitness function evaluation for IMRT beam angle optimization. *Cent. Eur. J. Oper.* **22**, 431–455 (2014). <https://doi.org/10.1007/s10100-013-0289-4>
3. Galvin, J.M., Chen, X., Smith, R.M.: Combining multileaf fields to modulate fluence distributions. *Int. J. Radiat. Oncol. Biol. Phys.* **27**, 697–705 (1993). [https://doi.org/10.1016/0360-3016\(93\)90399-g](https://doi.org/10.1016/0360-3016(93)90399-g)
4. Bortfeld, T.: IMRT: a review and preview. *Phys. Med. Biol.* **51**, R363–79 (2007). <https://doi.org/10.1088/0031-9155/51/13/R21>
5. Rocha H., Dias, J.M.: On the optimization of radiation therapy planning. Inesc Coimbra Research Reports (2009). [www.uc.pt/en/org/inescc/](http://www.uc.pt/en/org/inescc/)
6. Zarepisheh, M., Long, T., Li, N., Tian, Z., Romeijn, H.E., Jia, X., Jiang, S.B.: A DVH-guided IMRT optimization algorithm for automatic treatment planning and adaptive radiotherapy replanning. *Med. Phys.* **41**, 061711 (2014). <https://doi.org/10.1118/1.4875700>
7. Jia, Q., Li, Y., Wu, A., Guo, F., Qi, M., Mai, Y., Kong, F., Zhen, X., Zhou, L., Song, T.: OAR dose distribution prediction and gEUD based automatic treatment planning optimization for intensity modulated radiotherapy. *IEEE Access* **7**, 141426–141437 (2019). <https://doi.org/10.1109/ACCESS.2019.2942393>
8. Schipaanboord, B.W.K.M., Gizynska, K., Rossi, L., de Vries, K.C., Heijmen, B.J.M., Breedveld, S.: Fully automated treatment planning for MLC-based robotic radiotherapy. *Med. Phys.* **48**, 4139–4147 (2021). <https://doi.org/10.1002/mp.14993>
9. Bijman, R., Rossi, L., Janssen, T., de Ruiter, P., Carbaat, C., van Triest, B., Breedveld, S., Sonke, J.J., Heijmen, B.: First system for fully-automated multi-criterial treatment planning for a high-magnetic field MR-Linac applied to rectal cancer. *Acta. Oncol.* **59**, 926–932 (2020). <https://doi.org/10.1080/0284186X.2020.1766697>

10. Dias, J., Rocha, H., Ventura, T., Ferreira, B., Lopes, M.D.C.: Automated fluence map optimization based on fuzzy inference systems. *Med. Phys.* **43**, 1083–1095 (2016). <https://doi.org/10.1118/1.4941007>
11. Yan, D., Liang, J.: Expected treatment dose construction and adaptive inverse planning optimization: implementation for offline head and neck cancer adaptive radiotherapy. *Med. Phys.* **40**, 021719 (2013). <https://doi.org/10.1118/1.4788659>
12. Murthy, V., Master, Z., Adurkar, P., Mallick, I., Mahantshetty, U., Bakshi, G., Tongaonkar, H., Shrivastava, S.: “Plan of the day” adaptive radiotherapy for bladder cancer using helical tomotherapy. *Radiother. Oncol.* **99**, 55–60 (2011). <https://doi.org/10.1016/j.radonc.2011.01.027>
13. Barker, J.L., Garden, A.S., Ang, K.K., O’Daniel, J.C., Wang, H., Court, L.E., Morrison, W.H., Rosenthal, D.I., Chao, K.S., Tucker, S.L., Mohan, R., Dong, L.: Quantification of volumetric and geometric changes occurring during fractionated radiotherapy for head-and-neck cancer using an integrated CT/linear accelerator system. *Int. J. Radiat. Oncol. Biol. Phys.* **59**, 960–70 (2004). <https://doi.org/10.1016/j.ijrobp.2003.12.024>
14. Wieser, H.P., Cisternas, E., Wahl, N., Ulrich, S., Stadler, A., Mescher, H., Muller, L.R., Klinge, T., Gabrys, H., Burigo, L., Mairani, A., Ecker, S., Ackermann, B., Ellerbrock, M., Parodi, K., Jakel, O., Bangert, M.: Development of the open-source dose calculation and optimization toolkit matRad. *Med. Phys.* **44**, 2556–2568 (2017). <https://doi.org/10.1002/mp.12251>

# An Optimization Model for Power Transformer Maintenance



João Dionísio and João Pedro Pedroso

**Abstract** Power transformers are one of the main elements of a power grid, and their downtime impacts the entire network. Repairing their failures can be very costly, so sophisticated maintenance techniques are necessary. To attempt to solve this problem, we developed a mixed-integer nonlinear optimization model that, focusing on a single power transformer, both schedules this maintenance and also decides how much of the hourly demand it will satisfy. A high level of load on a power transformer increases its temperature, which increases its degradation, and so these two decisions have to be carefully balanced. We also consider that power transformers have several components that degrade differently. Our model becomes very difficult to solve even in reasonably sized instances, so we also present an iterative refinement heuristic.

**Keywords** Mathematical modeling · MINLP · Physics-Based optimization model · Scheduling · Production planning

## 1 Introduction

In this article, we present a MINLP with a convex relaxation for power transformer maintenance. We focus on accurately modeling the degradation process of its components, especially on the degradation of the insulating paper of the copper windings, which is where the majority of the nonlinearities reside. We also present two simple

---

J. Dionísio (✉) · J. P. Pedroso

Faculdade de Ciências, Universidade do Porto, rua do Campo Alegre s/n,  
4169-007 Porto, Portugal  
e-mail: [up201606210@up.pt](mailto:up201606210@up.pt)

J. P. Pedroso

e-mail: [jpp@fc.up.pt](mailto:jpp@fc.up.pt)

J. Dionísio

INESC TEC - Instituto de Engenharia de Sistemas e Computadores, Tecnologia e Ciência,  
Universidade do Porto, Porto, Portugal

J. P. Pedroso

CMUP, Centro de Matemática da Universidade do Porto, Porto, Portugal

© The Author(s), under exclusive license to Springer Nature Switzerland AG 2023

J. P. Almeida et al. (eds.), *Operational Research*, Springer Proceedings in Mathematics & Statistics 437, [https://doi.org/10.1007/978-3-031-46439-3\\_5](https://doi.org/10.1007/978-3-031-46439-3_5)

heuristics to try to get better solutions for larger instances, given that the model's difficulty increases steeply with the number of periods.

This article is organized as follows: Sect. 1 frames the problem and provides a light background on power transformers and maintenance practices, besides presenting the related work. In Sect. 2, the main model is presented, where we explain some constraints and the simplifications we have made. Here we also talk about the heuristics we developed. Section 3 presents the experimental setup and in Sect. 4 the results are shown. We conclude the paper in Sect. 5, where we give an overview of the article and talk about future research.

## ***1.1 Power Transformer***

A PT can be divided into subsystems, which are in turn composed of several components, each of which will suffer differently under the same conditions [2]. Furthermore, the condition of one component may influence the deterioration of another but not that of a third. Our choice of components is based on [1, 15]. It comprises of the cooling system (CS), the oil preservation system (OPS), the oil (O), and the insulating paper of the copper windings (W), this last one being a critical component that will put a hard limit on the transformer's remaining useful life (RUL, time until equipment failure under normal operation; see [17] for a literature review). The chosen components are essential to power transformer usage, interacting closely regarding maintenance and degradation, and were validated by experts in power transformer manufacturing.

## ***1.2 Maintenance Practices***

In maintaining power transformers, two maintenance strategies are very commonly used: Time-based Maintenance (TbM), where maintenance is planned considering the time since the last maintenance action, and Condition Based Maintenance (CbM), where the asset is regularly tested and when some predefined threshold is reached, maintenance is scheduled. Usually, most components are maintained with TbM in mind, but CbM has already seen some implementations as well [1]. A component associated with the latter strategy is the Oil, which is regularly tested to understand when maintenance should be performed. Some examples of CbM applied to PTs can also be found in the literature (see [7, 12]). The most common decision is then to combine TbM with CbM, depending on the considered components. However, both of these approaches make decisions using only locally available information, disregarding how the maintenance actions will impact the profit/cost/etc. over the long term. Considering that PTs are machines with a very long lifetime, it seems sensible and necessary to develop methods that take a more global view of this problem, which is the intent of this work.

### 1.3 Related Work

This problem can be seen as belonging to the class of production-maintenance scheduling problems with a single machine—see Geurtsen et al. wrote [8] for a comprehensive literature review. Many variants of this problem exist and have been studied for decades, albeit less than other more classical problems. Our work differs from the ones we could find in the sense that it is the first to research this topic in the context of power transformers specifically (although optimization models for PT maintenance have been developed, see the next paragraph). Additionally, nearly every paper looks at linear variants of this problem (see [9] for a resource-production exception), mostly because physical considerations tend to be ignored in more abstract models.

The scheduling of maintenance actions in PTs has been studied extensively [1, 4, 7, 10, 16], but to the best of our knowledge, it has never been attempted in conjunction with modeling the future condition of a PT. In these studies, the evolution of a PT's RUL tends to be overly simplistic, either modeled linearly or not modeled at all. Besides, most works in the literature focus on optimizing maintenance, ignoring that PTs can work under different loads, which has a large impact on their condition. Additionally, some of these papers focus on maintaining the PT in an optimal state [18], but this is not a realistic assumption.

Given that most models for PT maintenance tend to be linear, they can be solved to optimality in most reasonably sized instances, making heuristics not especially useful. The exception is Jahromi et al.'s work [5], a two-stage framework that divides the model into two parts, the first of which schedules maintenance on a long-term time horizon, and afterward, with that solution, operational considerations are taken. In this article, we present a heuristic with similar reasoning, but considerably different, given that it iteratively refines the maintenance decisions (rather than doing it one time), and operational decisions are always present.

## 2 Model

The objective of the model is to maximize the profit: revenue obtained by electricity sales with maintenance costs subtracted, where the electricity demand must not be exceeded, but also does not need to be satisfied. As for notation, variables will be denoted by lowercase letters (e.g.,  $r$ ) and parameters by capital letters (e.g.,  $R$ ), using different fonts. Table 1 details parameters and Table 2 details variables used.

Some of these parameters and variables will require additional superscripts. The  $C$  parameter, representing the cost associated with component maintenance, for example, will have different values for each component. As such, we will say  $C^W$ ,  $C^O$ ,  $C^{OPS}$  and  $C^{CS}$  to refer to the maintenance cost of the Winding, Oil, OPS, and CS components respectively. Other parameters will use the same notation. The variables and some parameters will also require two subscripts, detailing the time and the PT they are referring to. The RUL of the Winding component at time  $t$  is denoted by  $r_t^W$ .

**Table 1** Description of the parameters used in the model

| Parameter     | Parameter meaning                   | Parameter range                                         |
|---------------|-------------------------------------|---------------------------------------------------------|
| A             | Temperature's scaling factor        | $\mathbb{R}^+$                                          |
| B             | Temperature's growth factor         | $\mathbb{R}^+$                                          |
| C             | Cost of component replacement       | $\mathbb{R}^+$                                          |
| D             | Natural wear of component           | $[0, 1] \times \{\text{Oil}, \text{OPS}, \text{CS}\}$   |
| E             | Hourly electricity demand           | $[0, 2]$                                                |
| H             | Maximum permissible temperature     | $\mathbb{R}^+$                                          |
| Ht            | Hot-spot temperature reduction      | $\mathbb{R}^+$                                          |
| $\mathcal{K}$ | Components to be maintained         | $\{\text{Oil}, \text{OPS}, \text{CS}, \text{Winding}\}$ |
| L             | Effect of load in component wear    | $\mathbb{R}^+$                                          |
| M             | Big constant to dominate constraint | $\mathbb{R}^+$                                          |
| P             | Electricity price                   | $\mathbb{R}^+$                                          |
| Q             | Maximum permissible load            | $\mathbb{R}^+$                                          |
| R             | Maximum component RUL               | $\mathbb{R}^+$                                          |
| T             | Number of periods                   | $\mathbb{N}$                                            |
| V             | Moisture dependent degradation      | $\mathbb{R}^+$ [3, Table in page 21]                    |

**Table 2** Description of the variables used in the model

| Variable | Variable meaning                       | Variable type |
|----------|----------------------------------------|---------------|
| $h$      | Hotspot temperature, dependent on load | Continuous    |
| $m$      | Component maintenance                  | Binary        |
| $q$      | Transformer load                       | Continuous    |
| $r$      | Component remaining useful life        | Continuous    |
| $x$      | OPS RUL above 1/3                      | Binary        |
| $y$      | OPS RUL above 2/3                      | Binary        |

$$\underset{q, m}{\text{maximize}} \sum_{t=0}^T (P_t \cdot q_t - \sum_{k \in \mathcal{K}} C^k \cdot m_t^k) \quad (1a)$$

subject to

$$q_t \leq E_t, \forall t \in [0 .. T] \quad (1b)$$

$$r_0^k = R^k, \forall k \in \mathcal{K} \quad (1c)$$

$$r_t^O \leq r_{t-1}^O \cdot D^O - \frac{h_t}{H} + M_1 \cdot m_t^O, \forall t \in [1 .. T] \quad (1d)$$

$$r_t^{CS} \leq r_{t-1}^{CS} \cdot D^{CS} - L \cdot \frac{q_t}{Q} + M_2 \cdot m_t^{CS}, \forall t \in [1 .. T] \quad (1e)$$

$$r_t^{OPS} \leq r_{t-1}^{OPS} \cdot D^{OPS} - L \cdot \frac{q_t}{Q} + M_3 \cdot m_t^{OPS}, \forall t \in [1 .. T] \quad (1f)$$

$$r_t^W \leq r_{t-1}^W - 2^{\frac{h_t-98}{6}} + M_4 \cdot m_t^W, \forall t \in [1 .. T] \quad (1g)$$

$$r_t^W \leq r_{t-1}^W - V_y \cdot 2^{\frac{h_t-98}{6}} + M_4 \cdot m_t^W + y_t \cdot M_5 \cdot R^W, \forall t \in [1 .. T] \quad (1h)$$

$$r_t^W \leq r_{t-1}^W - V_x \cdot 2^{\frac{h_t-98}{6}} + M_4 \cdot m_t^W + x_t \cdot M_5 \cdot R^W, \forall t \in [1 .. T] \quad (1i)$$

$$q_t \leq Q \cdot \frac{r_t^k}{R^k}, \forall k \in \mathcal{K}, t \in [1 .. T] \in \text{PTs} \quad (1j)$$

$$m_t^{\text{OPS}} \leq m_t^O, \forall t \in [1 .. T] \quad (1k)$$

$$m_t^W \leq m_t^k, \forall k \in \mathcal{K}, t \in [1 .. T] \in \text{PTs} \quad (1l)$$

$$y_t \leq x_t, \forall t \in [1 .. T] \quad (1m)$$

$$3 \frac{r_t^{\text{OPS}}}{R^{\text{OPS}}} \geq y_t + x_t, \forall t \in [1 .. T] \quad (1n)$$

$$h_t \geq A \cdot e^{B \cdot q_t} - \text{Ht}^{\text{CS}} \cdot \frac{r_t^{\text{CS}}}{R^{\text{CS}}} - \text{Ht}^O \cdot \frac{r_t^O}{R^O}, \forall t \in [1 .. T] \quad (1o)$$

Constraints (1c)–(1i) are modeling the initial condition and its evolution for each of the components of the power transformer. A big-M term in Constraints (1d)–(1i) allows the modelling of the reduction in RUL (when the corresponding maintenance variable,  $m$ , is 0) and the replacement of the corresponding component (when  $m = 1$ ). The Winding component has three different on-off constraints due to the impact of moisture, whose constraints are present in (1m)–(1n).

With constraint (1j) we are conservatively limiting the load of the PT by the condition of its most damaged component, as a way to prevent unexpected (sometimes catastrophic) failure. In (1b), we are enforcing an upper bound based on the total demand, using consumption data from [14].

Constraints (1k) and (1l) are due to real-world limitations on power transformer maintenance, where some components need to be maintained when others are.

Constraint (1o) models the evolution of the PT temperature as a function of its load, as a simplification of the equations presented in [3]. Note that there is an exponential term here, that will be present in an exponent in constraints (1g)–(1i), conveying the very nonlinear degradation when the temperature raises above a certain threshold.

## 2.1 Simplifications

This model may quickly become very large, depending on the number of periods we are considering, but even for average-sized instances, it can be quite difficult to solve. As such, we were forced to concede some simplifications.

We assume that maintenance actions can only be taken on the first period of each year, in order to reduce the number of binary variables. The same is true for the indicator variables  $x$ ,  $y$ ,  $z$ .

Additionally, we use a representative day for each year. That is, every year has 24 periods, one per hour, that are to be interpreted as an indication of what is to be done

on every day of a particular year. We also measure the RUL of the components in 24h periods. The aim is again to reduce the size of the model, this time the number of continuous variables. As we are working with a MINLP, the continuous variables can be troublesome.

## 2.2 Iterative Refinement Heuristic

Without the aforementioned simplifications, the model would very quickly become intractable. To further aid in the solving process and to hopefully be able to lift some of these simplifications, we created an iterative refinement heuristic, which should reduce the computational burden of the (binary) maintenance variables.

In this heuristic, our objective is to decrease the impact of the binary variables in the solution time, by solving very coarse models when it comes to maintenance decisions and then refining the possibilities around the maintenance actions of the previous iteration. This idea is somewhat similar to what is presented by Jahromi et al. in [5], where a two-stage model first schedules maintenance on the long-term and then passes the decision to a medium-term scheduler. Our heuristic does not make this distinction between long and medium-term scheduling, always keeping a view of the global problem, and can also have more than one iteration.

---

### Algorithm 1: Iterative Refinement

---

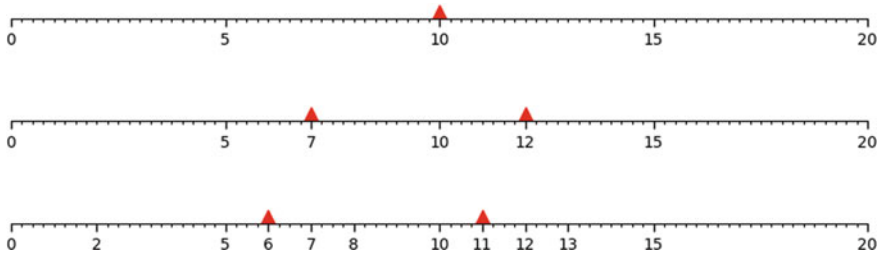
**Data:** Initial maintenance decisions, time/iteration/maintenance limit

```

1 while time, iterations, maintenance possibilities do not exceed limit do
2   build model with current maintenance possibilities;
3   if available, use previous best solution as starting point;
4   optimize model with given time limit;
5   if model is infeasible then
6     if maintenance possibilities exceed limit then
7       break;
8     double maintenance possibilities;
9     go to line 1;
10  if improvement is less than 1%, w.r.t. previous iteration then
11    solve to optimality;
12    break;
13  for possible maintenance variables do
14    if maintenance was scheduled then
15      add maintenance possibility at the midpoint between previous and current one;
16      add maintenance possibility at the midpoint between current and next one;
17 return best model;
```

---

The initial data consists of a set of initial coarse maintenance decisions, and an iteration and time limit. If these last two parameters are surpassed, we stop the algorithm (line 1) and return the best model found (line 17). Using the currently available maintenance possibilities, we build a model and optimize it (lines 2 and



**Fig. 1** Example of iterative refinement with a limit of 3 iterations

4), taking into account the remaining time we have, and using the best solution we have found if it exists (line 3). We know that any solution from an earlier iteration is feasible for further iterations. If the model only marginally improved (improvement of less than 1%, line 10), then we solve the model to optimality and return it (lines 11 and 17). This is a heuristic, so if it seems that the algorithm is converging to a local optimum, it exits before spending too much time. Assuming we did not exit the while loop, we iterate through each maintenance possibility, and if maintenance was scheduled (the corresponding binary variable equals 1), then we add two maintenance possibilities, in the midpoint between the current one and the previous, and in the midpoint between the current and the next one (lines 13 – 16).

Given high enough limits on time, iterations, and maintenance possibilities, this algorithm converges to a solution, as eventually, it will be unable to add new maintenance options.

Maybe counterintuitively, this method is not exact, as it may happen that the optimal maintenance configuration is impossible to reach in this manner. Consider an instance where the only profitable maintenance action consists of scheduling maintenance on year  $x$ , but year  $x$  is not a part of the initial coarse maintenance decisions. Even assuming we arrive at the optimal solution at every iteration, we will never add new maintenance possibilities, since maintenance is never scheduled.

Below, Fig. 1, exemplifies three iterations of this heuristic. First, maintenance is scheduled on year 10, which is then refined for year 15. On the final iteration, if the model can schedule maintenance on year 17, then it is profitable to do so in year 5 as well (due to the possibility of selling more electricity).

### 3 Benchmark Instances Used

In this section, we will detail the experimental setup that was used to validate the model, as well as the metrics used.

We will run the two different methods, the original model and the heuristic, for a varying number of years in the planning horizon, from 1 to 20, along with longer time periods, namely 50 and 100 years, all with a 2 hour time limit.

We run these experiments for 10 instances, that differ in their parameters, which were randomly sampled from what we understand to be sensible intervals. For the

two longest time frames, we also show the evolution of the incumbents of the two methods for a single instance. This instance's parameters correspond to the expected value of the aforementioned sampling.

The experiments were run on two 6-Core Intel Core i7 with 64 GB running at 3.20GHz. They ran the models with the SCIP solver, version 8.0.0 [6]. For the implementation, we used its Python interface, PySCIPOpt [13]. Some scripts (most importantly visualization with Matplotlib [11]) were implemented in Python version 3.7.9.

## 4 Results and Analysis

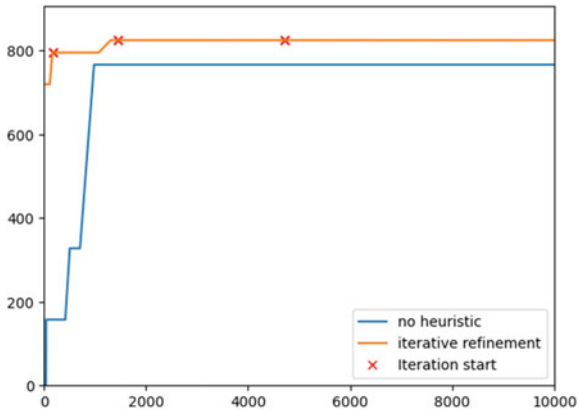
We were able to solve most instances to optimality in the given time limit, and those that we were not, had an acceptable gap, given their size. As for heuristics, iterative refinement did quite well for large instances, providing good incumbents, better than the original model, needing only a couple of iterations. Unfortunately, it only could reach optimal solutions (in relation to the original model), in very small instances.

Below we show the results obtained by the model (on the left) and by the heuristic (on the right). We decided not to show all the iterations of the heuristic because the resulting table would take up too much space.

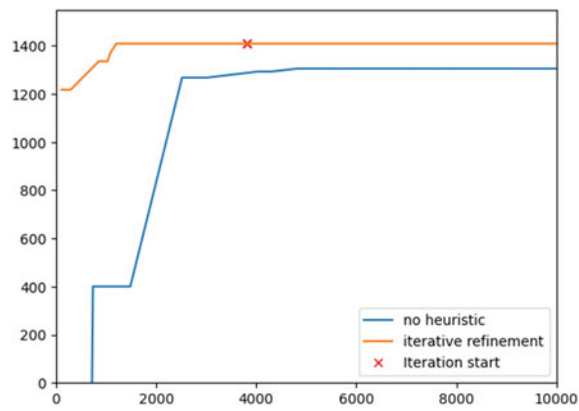
| Years | Incumbent | Time(s)  | Gap(%) | Years | Iterations | Incumbent | Time(s)   |
|-------|-----------|----------|--------|-------|------------|-----------|-----------|
| 1     | 28.8      | 0.15     | 0.0    | 1     | 1          | 28.8      | 0.09      |
| 2     | 56.1      | 0.32     | 0.0    | 2     | 1          | 56.1      | 0.22      |
| 3     | 82.4      | 0.68     | 0.0    | 3     | 1          | 82.4      | 0.418     |
| 4     | 107.8     | 0.97     | 0.0    | 4     | 1          | 107.8     | 0.569     |
| 5     | 130.9     | 1.59     | 0.0    | 5     | 1          | 130.9     | 0.844     |
| 6     | 152.0     | 3.95     | 0.0    | 6     | 1          | 151.2     | 1.194     |
| 7     | 175.8     | 3.4      | 0.0    | 7     | 1          | 168.9     | 1.219     |
| 8     | 199.2     | 6.31     | 0.0    | 8     | 1          | 184.0     | 1.688     |
| 9     | 221.1     | 23.5     | 0.0    | 9     | 1          | 196.2     | 2.5       |
| 10    | 241.5     | 23.27    | 0.0    | 10    | 1          | 205.6     | 2.624     |
| 11    | 259.6     | 9.64     | 0.0    | 11    | 4          | 259.6     | 12.741    |
| 12    | 277.4     | 18.22    | 0.0    | 12    | 4          | 277.4     | 15.858    |
| 13    | 293.2     | 19.38    | 0.0    | 13    | 4          | 291.5     | 17.188    |
| 14    | 306.0     | 56.88    | 0.0    | 14    | 4          | 305.6     | 28.479    |
| 15    | 319.8     | 80.32    | 0.0    | 15    | 4          | 319.8     | 30.182    |
| 16    | 332.6     | 122.0    | 0.0    | 16    | 4          | 332.4     | 51.78     |
| 17    | 344.2     | 165.45   | 0.0    | 17    | 4          | 344.2     | 49.954    |
| 18    | 354.9     | 390.02   | 0.0    | 18    | 4          | 354.8     | 102.473   |
| 19    | 370.8     | 822.02   | 0.0    | 19    | 4          | 364.5     | 152.982   |
| 20    | 390.5     | 785.71   | 0.0    | 20    | 4          | 373.3     | 125.074   |
| 50    | 766.1     | 10000.11 | 31     | 50    | 4          | 824.4     | 10000.016 |
| 100   | 1304.8    | 10000.0  | 49     | 100   | 2          | 1408.4    | 10000.033 |

The model solves all instances to optimality, except for the last two, where it obtains a relative gap of 31% and 49% respectively. The heuristic, on the other hand, performs much better than the original model in several instances. It is always much

**Fig. 2** Model versus  
Heuristic 50y



**Fig. 3** Model versus  
Heuristic 100y



faster (except for the latter two cases, where both methods reach the time limit), and it matches or surpasses the original model’s incumbent in 11/22, which is extended to 16/22 if we ignore differences of less than 1%. The heuristic surpasses the model precisely on the instances which the model cannot solve to optimality, for which we have below two graphs showing the evolution of the incumbent of the original model, in blue, and the iterative refinement heuristic, in orange (that we call “no heuristic” and “iterative refinement” respectively). Whenever a new iteration of the heuristic starts, we mark it with an “x” (Figs. 2 and 3).

In both of these cases, the heuristic managed to noticeably outperform the original model (7.61% and 7.94% improvement). We can also see that in these instances, subsequent iterations do not always produce better incumbents. Maybe this is the heuristic converging to a local optimum, or a case of starting to have too many maintenance decisions, thus diffculting the solving process.

Additionally, looking at the table results, we see that the gap is still significant, and the heuristic would most likely need additional rounds to arrive at the optimal solution, assuming it did not get stuck on a local optimum before that.

Analyzing the results, we can see that the heuristic underperforms more noticeably on instances 7–10, and then again on 19 and 20. Focusing on the first set of instances, looking back at Algorithm 1, we only create more maintenance possibilities if on the previous iteration a maintenance action was scheduled. Doing additional experiments showed us that in these instances, while the model performs maintenance actions, the heuristic does not. Without maintenance actions, the heuristic exits after a single iteration. This could be easily fixed, simply by adding an if condition that instead of exiting when no maintenance actions were scheduled, it runs the original model. This idea is a direct improvement over the heuristic we have now since only on relatively small instances would it be activated. More concretely, the only downside would be the time it takes to solve instances 7–10, but in the worst case (10 years), the heuristic would go from taking 2.6s, to taking  $2.6 + 23.3 = 25.9$ s, which is not at all unreasonable. We decided not to include this change because we think it illustrates well the importance of scheduling maintenance at the right time. In these cases, compared to the model’s optimal solution, if we can only schedule maintenance a couple of years before or after, then that maintenance action stops being profitable.

As for the second set of underperformances, the stopping criteria of the algorithm cannot explain the disparity. We believe that these instances start to become large enough to exhibit the behavior that we describe in Sect. 2.2, where we mention the possibility of this iterative refinement possibly missing the optimal scheduling.

As for recommendations, based on these results, we think that if the heuristic incorporates the changes we mentioned, then it should be preferred in nearly all cases, especially in very long time frames. However, if optimality is a requirement, then switching to the original model might be the only way to achieve it.

These results leave us very optimistic that we can decrease the granularity of small and medium-sized instances, and still obtain good solutions.

## 5 Conclusion and Future Work

In this article, we presented a MINLP for scheduling maintenance and operational decisions in a power transformer. To deal with the model’s difficulty, we came up with a heuristic that focuses on the difficulty presented by the number of binary variables, by starting with a subset of them and iteratively adding more.

We based the data characterizing a power transformer on the information we could discover in the literature [1], but we could not find satisfactory real-world data, which is desirable to validate the model.

The model almost exhibits a block angular structure, with a set of coupling variables (the load,  $q$  and the temperature,  $h$ ), and the blocks corresponding to the Winding + OPS and the Oil + Cooling System. The only constraints messing with the block structure are the maintenance dependency ones (1k)–(1l). Fortunately, maintenance actions are very rare, especially the ones concerning the OPS and the Winding (they are expensive), meaning that most “good” feasible solutions will satisfy these constraints. Preliminary experiments showed that fixing the load and the temperature

trivially solves the problem (as Constraint (1j) will essentially fix the maintenance as well), which indicates that Bender's decomposition may be a viable way to solve this problem.

**Acknowledgements** This work is financed by National Funds through the Portuguese funding agency, FCT—Fundação para a Ciência e a Tecnologia, within project LA/P/0063/2020.

## References

1. Guide for Transformer Maintenance. CIGRÉ, Paris (2011)
2. Guide on Transformer Intelligent Condition Monitoring (TICM) Systems. CIGRÉ, Paris (2015)
3. Power transformers—Part 7: Loading guide for mineral-oil-immersed power transformers. Standard, International Electrotechnical Commission (2018)
4. Abiri-Jahromi, A., Fotuhi-Firuzabad, M., Abbasi, E.: An efficient mixed-integer linear formulation for long-term overhead lines maintenance scheduling in power distribution systems. *IEEE Trans. Power Deliv.* **24**(4), 2043–2053 (2009)
5. Abiri-Jahromi, A., Parvania, M., Bouffard, F., Fotuhi-Firuzabad, M.: A two-stage framework for power transformer asset maintenance management part I: models and formulations. In: *IEEE Transactions on Power Systems*, vol. 28 pp. 1395–1403, July 2013
6. Bestuzheva, K., Besançon, M., Chen, W.-K., Chmiela, A., Donkiewicz, T., van Doornmalen, J., Eifler, L., Gaul, O., Gamrath, G., Gleixner, A., Gottwald, L., Graczyk, C., Halbig, K., Hoen, A., Hojny, C., van der Hulst, R., Koch, T., Lübbecke, M., Maher, S.J., Matter, F., Mühmer, E., Müller, B., Pfetsch, M.E., Rehfeldt, D., Schlein, S., Schlösser, F., Serrano, F., Shinano, Y., Sofranac, B., Turner, M., Vigerske, S., Wegscheider, F., Wellner, P., Weninger, D., Witzig, J.: The SCIP Optimization Suite 8.0. Technical report, Optimization Online, Dec. 2021. [http://www.optimization-online.org/DB\\_HTML/2021/12/8728.html](http://www.optimization-online.org/DB_HTML/2021/12/8728.html)
7. Dong, M., Zheng, H., Zhang, Y., Shi, K., Yao, S., Kou, X., Ding, G., Guo, L.: A novel maintenance decision making model of power transformers based on reliability and economy assessment. *IEEE Access* **7**, 28778–28790 (2019)
8. Geurtsen, M., Jeroen B.H.C., Didden, J.Ad., Atan, Z., Adan, I.: Production, maintenance and resource scheduling: a review. *European J. Oper. Res.* **305**(2), 501–529 (2023). <https://www.sciencedirect.com/science/article/pii/S0377221722002673>
9. Grigoriev, A., Uetz, M.: Scheduling jobs with time-resource tradeoff via nonlinear programming. *Disc. Optim.* **6**(4), 414–419 (2009). <https://www.sciencedirect.com/science/article/pii/S1572528609000334>
10. Guo-Hua, Q., Zheng, R., Lei, S., Bo, Z., Jian-Gang, X., Xiang-Ling, Z.: A new life cycle cost model of power transformer and its comprehensive sensitivity analysis. In: 2014 International Conference on Power System Technology, pp. 1342–1348 (2014)
11. Hunter, J.D.: Matplotlib: a 2d graphics environment. *Comput. in Sci. Eng.* **9**(3), 90–95 (2007)
12. Liang, Z., Parlikad, A.: A markovian model for power transformer maintenance. *Int. J. Electr. Power Energy Syst.* **99**, 175–182 (2018). <https://www.sciencedirect.com/science/article/pii/S0142061517321312>
13. Maher, S., Miltenberger, M., Pedroso, J.P., Rehfeldt, D., Schwarz, R., Serrano, F.: PySCIOpt: Mathematical programming in python with the SCIP optimization suite. In: *Mathematical Software—ICMS 2016*, pp. 301–307. Springer International Publishing (2016)
14. Palmer, J., Cooper, I.: United Kingdom Housing Energy Fact File, Mar. 2014
15. Reclamation, B.: Transformers: Basics, Maintenance and Diagnostics. US Department of the Interior Bureau of Reclamation, Denver, Colorado, USA (2005)

16. Sarajcev, P., Jakus, D., Vasilj, J.: Optimal scheduling of power transformers preventive maintenance with bayesian statistical learning and influence diagrams. *J. Clean. Prod.* **258**, 120850 (2020). <https://www.sciencedirect.com/science/article/pii/S0959652620308970>
17. Si, X.-S., Wang, W., Hu, C.-H., Zhou, D.-H.: Remaining useful life estimation—a review on the statistical data driven approaches. *European J. Oper. Res.* **213**(1), 1–14 (2011). <https://www.sciencedirect.com/science/article/pii/S0377221710007903>
18. Vahidi, B., Zeinoddini-Meymand, H.: Health index calculation for power transformers using technical and economical parameters. *IET Sci. Meas. Technol.* **10**, 07 (2016)

# Multi-objective Finite-Domain Constraint-Based Forest Management



Eduardo Eloy, Vladimir Bushenkov, and Salvador Abreu

**Abstract** This paper describes an implementation of a Constraint Programming approach to the problem of multi-criteria forest management optimization. The goal is to decide when to harvest each forest unit while striving to optimize several criteria under spatial restrictions. With a large number of management units, the optimization problem becomes computationally intractable. We propose an approach for deriving a set of efficient solutions for the entire region. The proposed methodology was tested for Vale do Sousa region in the North of Portugal.

**Keywords** Forest management · Constraint programming · Constraint modeling · Multi-criteria optimization · Pareto frontier

## 1 Introduction

This paper presents an extension to previous work [1] of the authors on single criteria Forest Management optimization by considering multiple-criteria. Forest Management remains an activity of prime ecological importance where the interests of multiple stakeholders can lead to complex combinatorial and optimization problems, more so when different measures of economic performance such as wood yield and cash flow have to be balanced with environmental impact measures such as soil loss and fire resistance.

---

E. Eloy (✉)  
University of Évora, Evora, Portugal  
e-mail: [eeloy@di.uevora.pt](mailto:eeloy@di.uevora.pt)

V. Bushenkov  
CIMA, University of Évora, Evora, Portugal  
e-mail: [bushen@uevora.pt](mailto:bushen@uevora.pt)

S. Abreu  
NOVA-LINCS, University of Évora, Evora, Portugal  
e-mail: [spa@uevora.pt](mailto:spa@uevora.pt)

**Fig. 1** Vale De Sousa forest and its regions



Traditionally these problems are modelled and implemented with Integer Programming approaches and adjacency constraints [2–5], which work by applying conditions to when activities like harvesting can be applied to adjacent units of land. These approaches offer varying results depending on the model and the specific forest case. This paper describes an approach utilizing Constraint Programming (CP) [6] to solve a concrete problem of multiple-objective optimization, those objectives being wood yield, soil loss and fire resistance.

The effort to develop this approach was made possible by the MODFIRE project, which provided us with funding as well as all the relevant data regarding the forest which we based this implementation on, the Vale de Sousa forest, with the goal of optimizing multiple objectives when managing the forest throughout the 50 year planning horizon (2020–2070).

The forest is divided into 1406 Management Units (MUs), as shown in Fig. 1.

For each MU several different *prescriptions* were defined. Each prescription provides an option as to when each MU should be harvested (cut down every tree and making a clear-cut) or its branches thinned, as well as what “reward” is obtained by applying either of those actions in terms of the different optimizable criteria. It is assumed that in each MU there is only 1 species of tree which is specified by the prescriptions, with the possible tree species being eucalyptus (Ec), cork oak (Sb), pine tree (Pb), chestnut tree (Ct), pedunculate oak (Qr) and different riparian species of trees (Rp).

The main restriction regarding the harvest action is that to prevent issues of soil erosion the environmental authority imposes a limit to the *continuous amount of forest area that can be harvested (cut down)* and this limit is generally set at 50ha.

The combinatorial problem is thus as follows: to attribute a prescription to each MU so that in any given year, the continuous adjacent area of MUs where the action to be taken that year is a harvest must be below 50ha, meaning there cannot be very large sections of continuous forest area being deforested in any single year. To simplify the model, restrictions on wood flow volume were not considered.

The paper is structured as follows: after introducing the previous problem statement, we briefly review the state of the art in sect. 2 and then succinctly recap Constraint Programming (CP) in Sect. 3. In Sect. 4 we discuss the implementation and computational environment and evaluate the performance of the system in Sect. 5. Section 6 concludes this paper with a brief analysis and possible directions for further development.

## 2 Related Work

As previously stated the present work comes as a continuation of our participation in the BIOECOSYS project, as documented in [1], and of our current participation in the MODFIRE project. In that project we worked with a different structure for the Vale de Sousa Forest: the Paiva sub-region was divided into north and south and Penafiel included a cluster of several large MUs which made it a hard sub-region to solve. This was because a valid combination of prescriptions where none of the MUs composing these clusters were harvested at the same time was never found by the solver. Moreover, the planning horizon was 90 years instead of 50 years and the prescriptions only came associated with one optimizable criterion which was the wood yield.

The implementation we have now is more complete and able to also model and deal with multi-criteria optimization.

In other methods and models for forest planning problems, such as the ones cited in this paper's introduction, the authors usually represent the forests as adjacency graphs, so we did as well.

In a 1999 paper Alan T. Murray [7] proposed modeling constraints on the maximum harvest area by allowing multiple adjacent units to be harvested as long as their combined area does not exceed a pre-defined limit. We adopt this approach, termed Area Restriction Model (ARM), with a parametric area limit.

Latter Constantino et al. [8] proposed for this problem a compact mixed integer linear programming model, with polynomial number of variables and constraints. Multi-criteria and bi-level approaches have also been applied in forest management [9–11].

### 3 Constraint Programming

A good way to understand CP is that it stands as a synthesis of Mathematical Programming, where the programmer specifies the problem and the system attempts to find a solution on its own, and Computer Programming, where the programmer has to instruct the system on how to go about finding a solution. In CP the programmer specifies the problem as a set of relations that must hold (not be broken) but they may also give hints on how the system could go about finding a solution.

An application may be formulated as a *Constraint Satisfaction Problem* (CSP)  $\mathcal{P}$ , which consists of a triple  $(V, D, C)$  where  $V$  is a set of *variables*,  $D$  is a set of *domains* for the elements of  $V$  and  $C$  is a set of *constraints*, i.e. relations over  $\mathcal{P}(D)$  which must hold. The nature of the domains for the variables (*Finite Domains*, Booleans, Sets, Real numbers, Graphs, etc.), together with the specific relations (i.e. the *Constraints*) greatly influence the class of problems and application areas for which Constraint Programming form a good match.

A *Constraint Optimisation Problem* (COP) is like a CSP but we are also interested in minimizing (or maximizing) an *objective function*. To achieve this, one may equate the objective function to the value of a particular variable. It is then possible to solve a COP by iteratively solving interrelated CSPs, involving the addition of a constraint which establishes an inequation between the analytical definition of the objective function and the previously found value.

The model for an application problem may be declaratively formulated as a CSP, which will form the specification for a *constraint solver* to find a solution thereto. Many successful approaches have been followed to solve CSPs, namely systematic search, in which variables see their domain progressively restricted and each such step triggers the reduction of the domains of related variables, as dictated by the *consistency* policy—these are in general designated as propagation-based constraint solvers and there are several ones, some being presented as libraries for use within a general-purpose programming language, such as Gecode [12] or Choco [13]. Others offer a domain-specific language (DSL) which may be used to model a problem and provide it as input to different solvers; such is the case for instance for MiniZinc [14] or PyCSP3 [15].

Another approach entails selecting an initial solution candidate and working a path towards an actual solution by means of an *iterative repair* algorithm. The latter forms the basis for several *local search* techniques, which may be generalised to related methods called *metaheuristics*. Solvers which derive the strategy used to guide the search from the specification of a CSP are called *constraint-based local search* solvers [16].

Solvers exist for both propagation-based search and metaheuristic search, which exhibit high performance and the capacity to make use of parallel hardware to attain yet better performance, e.g. as discussed in [17, 18].

Constraint modeling allows one to express a problem by means of both simple arithmetic and logic relations, but also resorting to *global constraints*. These are instance-independent yet problem-class-specific relations, for which particular ded-

icated algorithms can be devised and encapsulated in a reusable specification component. Intuitively, a global constraint expresses a useful and generic higher-level concept, for which there is an efficient (possibly black-box) implementation. For instance, the `AllDifferent` constraint applies to a set of variables and requires them to take pairwise distinct values. It may be internally implemented in a naïve way by saying that each distinct pair of variables in the list must be different, or it may resort to a more specialized algorithm to achieve the same result more efficiently. The application programmer will benefit from the performance gain with no additional effort.

Global constraints have proved to be a fertile ground for effective research, over the years. A limited common set of global constraints has been presented in the XCSP<sup>3</sup>-core document [19], which lists 20 such frequently used and generally useful constraints. This forms the basic vocabulary of XCSP<sup>3</sup>-core, an intermediate representation for CSPs designed with the purpose of interfacing different high-level modeling tools with distinct specific constraint solvers.

## 4 Implementation

After experimenting with other tools of CP we settled on the Choco Solver [13], an open-source Java library for CP. This tool allowed us to mix the versatility of the java language with a suite of known constraints that come implemented with Choco such as the “Sum Constraint” where the sum total of values in an array is constrained to be within a set range of values or to be equal to a value. But crucially Choco allows for custom constraints to be created by defining how the effect of said constraints are propagated, this turns out to be how we create and apply the constraints needed to enforce the 50ha limit of continuously harvested forest area. We opted to base our work on the Choco Java constraint programming library, as we had already developed a dedicated global constraint in the form of a custom propagator. Other constraint programming libraries which we considered using include Gecode [20] and IBM ILOG CP Optimizer [21].

### 4.1 Model Description

The process of implementing the problem entails firstly setting up 2 data structures, one is an array of `Node` objects called `Nodes` where each `Node` contains the area of the MU and the IDs of its adjacent MUs. The other array is called `MUS`, it has as many variables as there are MUs in the forest, it is a constraint variable array so each variable does not have a set value but a domain of possible values. These possible values are the possible prescriptions that may be applied in the corresponding MU and, after setting up the constraints, during the solving process the solver will attempt remove values from the domains of these variables as it attempts to find a state where

the every variable has a value assigned to it that does not break the constraints between the variables. Another constraint variable array is called `WoodYields`, each index of the array corresponds to the wood yielded by harvesting/thinning one of the MUs in the forest when applying the prescription that the solver picked for the MU (this process involves using the “Element” constraint from the Choco-Solver framework). This is done so that once a solution is found the contents of this array can be summed up to obtain the total amount of wood yielded by that solution. This is done analogously for the other criteria we want to optimize. Once the setup is done we iterate through the main loop, as shown in Listing 1.

**Listing 1** Main Loop

```
for (Var node in Nodes) {
    CreateConstraint (node, MUS, MAXLIMIT);
}
SumOfWood = sumConstraint(WoodYields, range=[-999999, 999999]);
SumOfCrit2 = sumConstraint(Crit2Totals, range=[-999999, 999999]);
```

This loop iterates through every MU in the input and imposes all valid constraints pertaining to it and its possible prescriptions. These constraints are implemented as a global constraint, via a custom propagator, as shown in Listing 2.

The propagator essentially iterates through the MU’s possible prescriptions and checks if a prescription value can be applied by recursively checking its neighbouring MUs and their possible prescriptions. If at any one point the total sum of continuous forest area cut down exceeds the given limit, the propagator fails and another value will be chosen. The propagator calls a recursive function which verifies that a given MU is valid, with reference to the maximum cut area requirement, as shown in Listing 3.

**Listing 2** Custom Propagator

```
void propagate (){
    for (int year in node.yearsWithCuts) {
        try {gladePropagate(node, year, 0)}
        catch (Exception limitSurpassed) { fails(); }
    }
}
```

**Listing 3** Propagator Helper

```
int gladePropagate (node, year, sum) {
    if (node.hasCut() && node.isValid()) {
        sum += node.area;
        if (sum > MAXLIMIT) { throw Exception; }
        for (neighbourNode in node.neighbours())
            sum = gladePropagate (neighbourNode, year, sum);
    }
    return sum;
}
```

After leaving the main loop, the model has been fully setup.

If the problem is one of constraint satisfaction and not optimization the solver can now be activated and the solution, if found, will be written onto an input file.

The default search strategy that Choco-Solver uses to assign integer values to the constraint variables is based on attributing weights to these values and starting the assignment process with the lowest bound variable so the results are deterministic.

## **4.2 Multi-criteria Optimization**

Multi-Criteria optimization requires the user to inform the solver of which criteria to optimize, and the solver cannot mix and match maximization problems with minimization problems. If the user wants to maximize most variables and minimize others, he must convert the minimization problems into maximization problems by switching the sign of their criteria, as was done in the case of the Vale de Sousa forest with soil loss.

Then, the solver runs on a loop to find as many valid solutions as possible, so solutions that simply do not break the 50ha limit.

Theoretically, it is going to reach every possible valid solution. In practise, the search stops with a memory error or, given that the search may take too long otherwise, when a time limit set by the user is reached.

Once the loop terminates the solver utilizes the objective criteria of all the solutions it found as points to calculate the Pareto efficient points and therefore the Pareto efficient solutions, which it writes to an output file so we can plot the Pareto Frontier and pick a point to paint on an output map.

## **5 Experimental Evaluation**

The testing was done on a laptop running Ubuntu 20.04.3 LTS, with 4GB of available RAM and 4 cores. The code was compiled using the Java SDK version 8.

It should be noted that no solution for the full problem could be found in a reasonable time frame on this platform: the complete problem includes 1406 management units and the program ran for an entire day and a solution was not found. Consequently we opted for an approximation to the problem.

## 5.1 *Solving the Problem by Sub-region*

### **Paredes and Penafiel**

In regards to simple constraint satisfaction problems valid solutions are always found for both of these sub-regions separately and even when solved together.

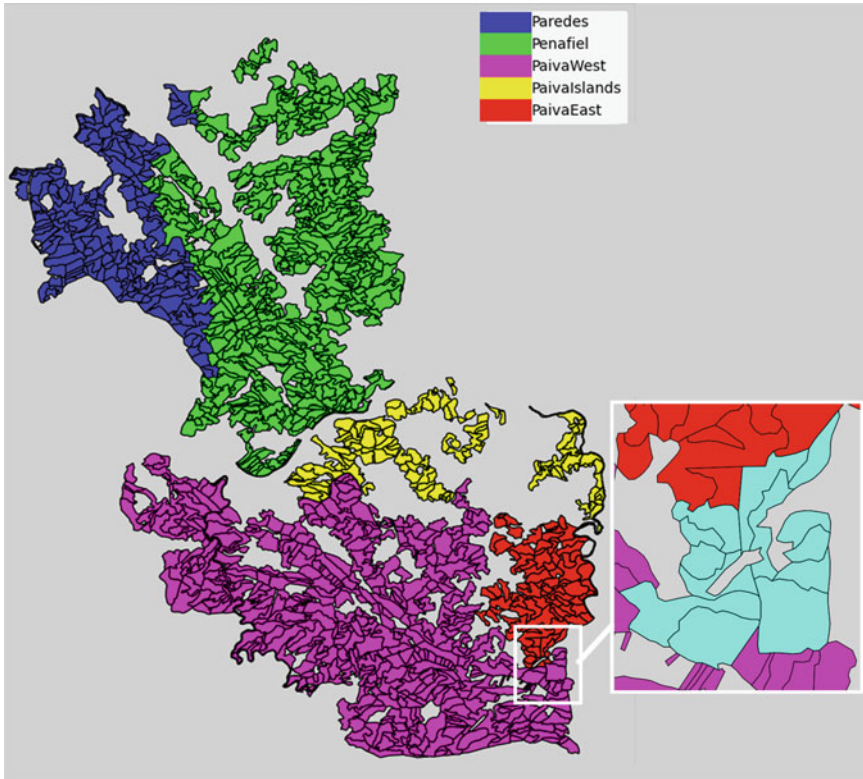
Regarding multi-criteria optimization for Paredes and Penafiel together the results depend of the number of criterion which we are trying to simultaneously optimize, this type of optimization depends on the solver finding multiple valid solutions. There is a complexity cost in the number of constraints for each criterion the solver has to consider when finding these solutions and while it generally will find valid solutions when optimizing 1 to 4 criteria, depending on the criterion, more than that will leave the solver running for a long time without finding a valid solution.

Being able to find valid solutions for Paredes and Penafiel together is significant since these 2 sub-regions are completely separated from the Paiva sub-region so they can be treated as a separate problem from Paiva.

### **Paiva**

The Paiva sub-region is a source of difficulty for the implementation because the solver is left running for hours and hours and a valid solution is never be found. As this sub-region proved to be the most complex in the forest we deemed it expedient to divide the Paiva sub-region into 3 separate sub-regions which are individually solvable, with the goal of building a Pareto frontier for each of the sub-regions and “joining” the Pareto frontiers for a complete Paiva Pareto frontier. By closely observing the sub-region an attempt was made to divide Paiva in spots where there are few adjacent MUs between sub-regions, ending up with PaivaWest (Pink), PaivaEast (Red) and PaivaIslands (Yellow). PaivaWest is completely separated from PaivaIslands, PaivaEast is only connected to PaivaIslands through 1 MU and PaivaWest is connected to PaivaEast through few MUs. So the forest is divided as shown in Fig. 2.

Because this “joining” of the solutions is being done outside of the Constraint Programming implementation, there is no guarantee that the solutions in this complete Paiva Pareto frontier will not break the 50ha area limit. Therefore we decided to first find a single valid solution to the “Contact Zone” (the group of MUs that are close to the border of 2 adjacent sub-regions) between PaivaWest and PaivaEast, also shown in Fig. 2, and “lock” the pairs of MU/Prescriptions of that solution for the solving process of PaivaEast and West, guaranteeing that the solutions found enforce the 50ha limit between sub-regions. We also developed a script to verify if a solution breaks the 50ha limit.



**Fig. 2** Vale de Sousa forest divided by sub-region and contact zone

## 5.2 *Joining the Sub-regions*

Let's assume that we have 2 sub-regions with  $N_1$  and  $N_2$  Pareto points respectively. This “joining” process works by “adding” each Pareto efficient point of one sub-region to each Pareto efficient point of the other sub-region, yielding  $N_1 \times N_2$  points. Since this may result in a large number of points some filtration is required to reduce this number. At first we round up the objective values of each point (for example by 4 decimal digits) and then we eliminate the non-efficient points (as well as points with the same objective values). After that we combine the “filtered” points of one sub-region with the “filtered” points of the other sub-region. From the set of obtained points we remove the dominated ones.

Other sub-regions are “joined” with the resulting Pareto set in the same manner successively.

Once the final Pareto frontier is built the user can pick a point on the Decision Map (see Fig. 3), learn the effects of applying the corresponding solution and observe its application throughout the planning horizon.

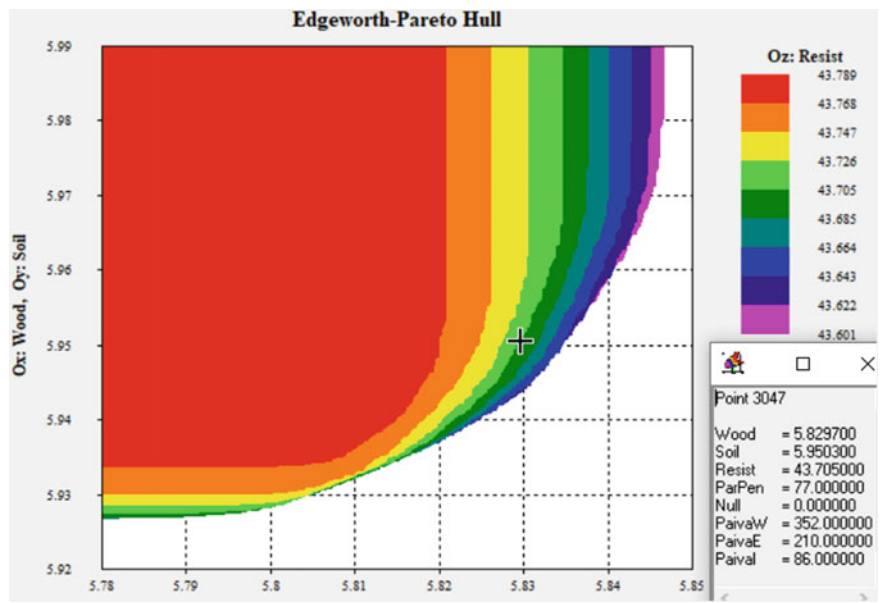


Fig. 3 Visualization of 14833 Pareto points for the complete Vale de Sousa forest

5.3 Results

Firstly valid solutions were found for 3 criteria, wood yield, soil loss and fire resistance, for all sub-regions. The solver was left running until a memory error occurred and all the solutions found until that point were considered, then the Pareto efficient solutions of each sub-region were joined as previously explained. The execution time until a memory error occurs varies depending on the sub-region, as show in Table 1

In our implementation to “join” the results firstly the Pareto points of PaivaWest are joined with the ones from PaivaEast, yielding 103246 possible points but after filtering out the dominated points only 3498 remained. Then once these results are joined with PaivaIslands 9166 points are obtained. Finally, joining these points with

Table 1 Execution details with all sub-regions

| Sub-region     | Execution time (h) | Valid solutions | Pareto solutions | After filtering |
|----------------|--------------------|-----------------|------------------|-----------------|
| ParPen         | 04:08              | 9932            | 813              | 226             |
| PaivaIslands   | 06:00              | 21555           | 578              | 284             |
| PaivaWest      | 00:46              | 9364            | 742              | 247             |
| PaivaEast      | 03:40              | 11534           | 874              | 418             |
| CompleteForest | –                  | –               | 16901            | 14833           |

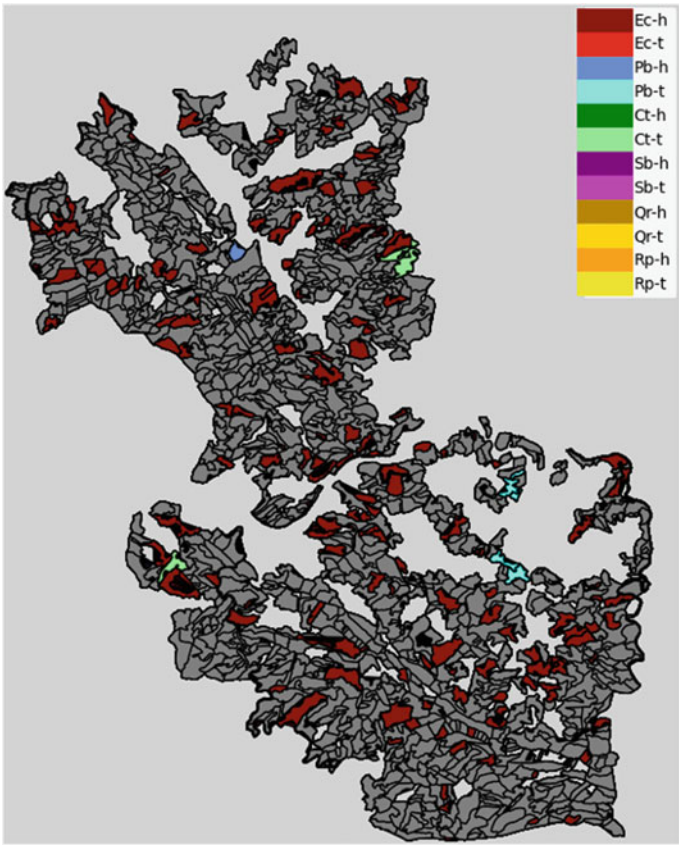
the 226 filtered points from Paredes+Penafiel yields 16901 points (or complete forest solutions), furthered reduced to 14833 points after filtering.

A point is then picked utilizing a visualizer [22, 23], this visualization is obtained by drawing the dominance cones of each point in the Pareto frontier, as shown in Fig. 3 the vertices of these cones are distributed uniformly so we may say that in this case we obtained well distributed points.

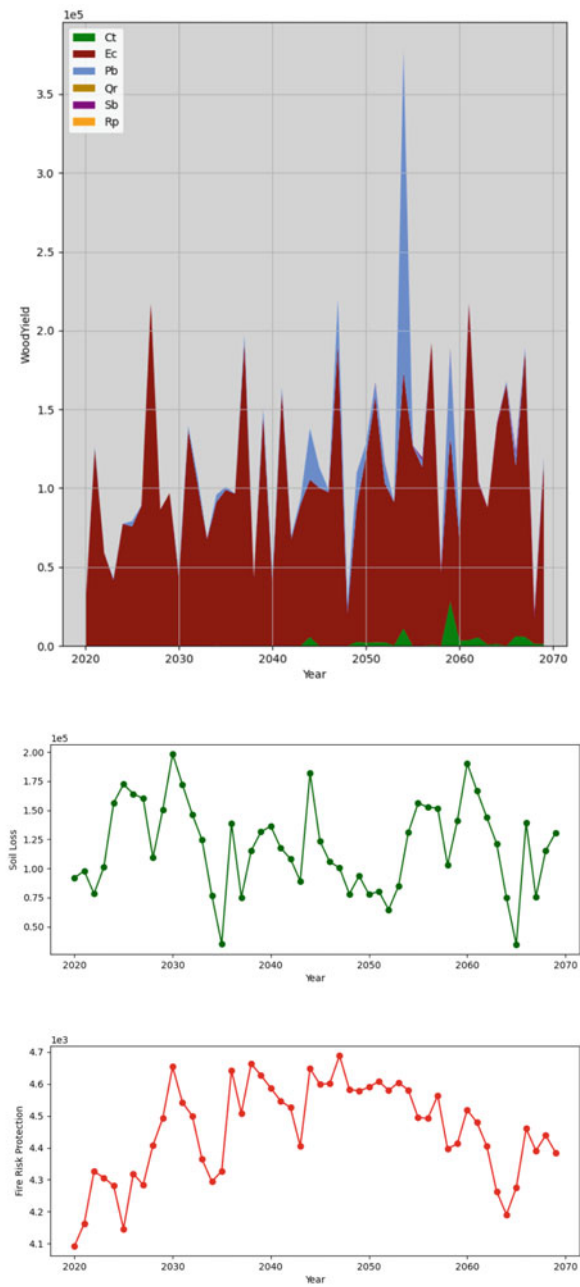
The maximum wood yield represented in Fig. 3 is around 5.845 tons of wood, the minimum soil loss is 5.928 tons of soil lost and the maximum fire resistance index is 43.789.

To showcase the results we opted to use as an example the compromise point marked by a cross in Fig. 3, yielding 5.8297 tons of wood, 5.9503 tons of soil lost and 43.705 fire resistance index.

The point we picked as an example has a corresponding solution which assign a prescription number to each MU. This solution may then be visualized on maps rep-



**Fig. 4** Example of an output map for a full solution of the Vale de Sousa forest, year 2041



**Fig. 5** Graphs showing the year by year values of the wood yield, soil loss and fire risk protection of the chosen solution

representing management units on which the harvest or thinning operations are applied in each year throughout the 50 year planning horizon.

Figure 4 showcases an output map corresponding to the year 2041 where the colors represent both the species of tree planted in the MUs and the action to be taken upon them, for example dark red represents harvest of eucalyptus (Ec-h) and light green represents thinning of chestnut trees (Ct-t). Figure 5 showcases 3 graphs visualizing the values of the criteria obtained each year.

## 6 Conclusions and Future Work

In this work we solved the large scale multi-objective Vale de Sousa forest managing problem with adjacency constraints by applying a Constraint Programming approach.

The problem was divided into computationally manageable sub-problems and for each of them a Pareto frontier was calculated by using the Choco-Solver.

These partial Pareto frontiers were combined into a Pareto frontier for the entire Vale de Sousa region. The proposed methodology allowed us to obtain solutions for large scale multi-objective forest scheduling problems in a bounded time frame.

In future developments, we'll study the Vale de Sousa problem with more than 3 criteria being considered and work on improving the constraint programming model and using metaheuristic search procedures.

**Acknowledgements** The authors acknowledge the Portuguese Fundação para a Ciência e a Tecnologia (FCT) for supporting the development of the system presented in here through the MODFIRE grant (PCIF/MOS/0217/2017).

## References

1. Eloy, E., Bushenkov, V., Abreu, S.: Constraint modeling for forest management. In: Tchemisova, T.V., Torres, D.F.M., Plakhov, A.Y. (eds.), *Dynamic Control and Optimization*, Cham, pp. 185–200. Springer International Publishing (2022)
2. Hof, J., Joyce, L.: A mixed integer linear programming approach for spatially optimizing wildlife and timber in managed forest ecosystems. *Forest Sci.* **39**, 816–834 (1993)
3. McDill, M.E., Rebain, S., Braze, M.E., McDill, J., Braze, J.: Harvest scheduling with area-based adjacency constraints. *Forest Sci.* **48**(4), 631–642 (2002)
4. Gunn, E., Richards, E.: Solving the adjacency problem with stand-centered constraints. *Canadian J. Forest Res.* **35**, 832–842 (2005)
5. Gharbi, C., Ronnqvist, M., Beaudoin, D., Carle, M.-A.: A new mixed-integer programming model for spatial forest planning. *Canadian J. Forest Res.* **49**, 1493–1503 (2019)
6. Apt, K.: *Principles of Constraint Programming*. Cambridge University Press (2003)
7. Murray, A.: Spatial restrictions in harvest scheduling. *Forest Sci.* **45**(1), 45–52 (1999)
8. Constantino, M., Martins, I., Borges, J.G.: A new mixed-integer programming model for harvest scheduling subject to maximum area restrictions. *Oper. Res.* **56**(3), 542–551 (2008)
9. Pascual, A.: Multi-objective forest planning at tree-level combining mixed integer programming and airborne laser scanning. *Forest Ecol. Manag.* **483**, 118714 (2021)

10. Marques, S., Bushenkov, V., Lotov, A., Borges, J.G.: Building pareto frontiers for ecosystem services tradeoff analysis in forest management planning integer programs. *Forests* **12**(9) (2021)
11. Marques, S., Bushenkov, V.A., Lotov, A.V., Marto, M., Borges, J.G.: Bi-Level participatory forest management planning supported by pareto frontier visualization. *Forest Sci.* **66**, 490–500 (2019)
12. Schulte, C., Tack, G., Lagerkvist, M.Z.: Modeling. In: *Modeling and Programming with Gecode*, Christian Schulte and Guido Tack and Mikael Z. Lagerkvist. Corresponds to Gecode 6.2.0 (2019)
13. Prud'homme, C., Fages, J.-G., Lorca, X.: Choco Solver Documentation. TASC, INRIA Rennes, LINA CNRS UMR 6241, COSLING S.A.S. (2016)
14. Nethercote, N., Stuckey, P.J., Becket, R., Brand, S., Duck, G.J., Tack, G.: Minizinc: towards a standard CP modelling language. In: Bessiere, C. (ed.), *Principles and Practice of Constraint Programming—CP 2007*, 13th International Conference, CP 2007, Providence, RI, USA, 23–27 Sept. 2007, Proceedings, vol. 4741 of *Lecture Notes in Computer Science*, pp. 529–543, Springer, 2007
15. Lecoutre, C., Szczepanski, N.: Pycsp3: Modeling combinatorial constrained problems in python (2020). [arXiv:2009.00326](https://arxiv.org/abs/2009.00326)
16. Michel, L., Hentenryck, P.V.: Constraint-based local search. In: Martí, R., Pardalos, P.M., Resende, M.G.C. (eds.), *Handbook of Heuristics*, pp. 223–260. Springer (2018)
17. Codognet, P., Munera, D., Diaz, D., Abreu, S.: Parallel local search. In: Hamadi, Y., Sais, L. (eds.), *Handbook of Parallel Constraint Reasoning*, pp. 381–417. Springer (2018)
18. Régim, J., Malapert, A.: Parallel constraint programming. In: Hamadi, Y., Sais, L. (eds.) *Handbook of Parallel Constraint Reasoning*, pp. 337–379. Springer (2018)
19. Boussemart, F., Lecoutre, C., Audemard, G., Piette, C.: Xcsp3-core: a format for representing constraint satisfaction/optimization problems (2020). [arXiv:2009.00514](https://arxiv.org/abs/2009.00514)
20. Schulte, C., Stuckey, P.J.: Efficient constraint propagation engines. *ACM Trans. Program. Lang. Syst.* **31**(1), 2:1–2:43 (2008)
21. Laborie, P., Rogerie, J., Shaw, P., Vilím, P.: IBM ILOG CP optimizer for scheduling. *Constraints* **23**, 210–250 (2018). Apr
22. Lotov, A., Kamenev, G., Berezkin, V.: Approximation and visualization of pareto-efficient frontier for nonconvex multiobjective problems. *Dokl* **66**, 260–262 (2002)
23. Lotov, A., Bushenkov, V., Kamenev, G.: Interactive decision maps: approximation and visualization of pareto frontier. Springer (2004)

# Implementation of Geographic Diversity in Resilient Telecommunication Networks



Maria Teresa Godinho and Marta Pascoal

**Abstract** In the  $K$   $D$ -geodiverse shortest paths problem we intend to determine  $K$  paths with the minimum total cost between two nodes of a network, such that the distance between the internal elements of each pair of the  $K$  paths exceeds a given value  $D > 0$ . Internal elements of a path are all the nodes and arcs that constitute it, except the first and the last nodes as well as the arcs incident on those nodes. This problem models disaster situations that affect a geographic zone with an impact radius  $D$ , guaranteeing the connection between the first and the last nodes, provided that such an event affects, at most,  $K - 1$  of these paths. In particular, geodiverse protection mechanisms are used to ensure the resilience of telecommunications networks. We extend the application of geodiversification found in past literature for  $K = 2$ , and find  $K > 2$  paths between a source and a destination nodes with the minimum total cost for a given distance value  $D$ , through an Integer Linear Programming formulation. Computational results of the proposed model are presented when applied to two reference networks in telecommunications.

**Keywords** Telecommunications routing · Geodiversity · Integer linear programming

---

M. T. Godinho (✉)

Instituto Politécnico de Beja, Department of Mathematical and Physical Sciences, 7800-295 Beja, Portugal

Cmaf-cIO, Universidade de Lisboa, 1749-016 Lisboa, Portugal

e-mail: [mtgodinho@ipbeja.pt](mailto:mtgodinho@ipbeja.pt)

M. Pascoal

University of Coimbra, CMUC, Department of Mathematics, 3001-501 Coimbra, Portugal

e-mail: [marta@mat.uc.pt](mailto:marta@mat.uc.pt)

Institute for Systems Engineering and Computers – Coimbra, rua Sílvio Lima, Pólo II, 3030-290 Coimbra, Portugal

Dipartimento di Elettronica, Informazione e Bioingegneria, Politecnico di Milano, Piazza Leonardo da Vinci, 32, Milan 20133, Italy

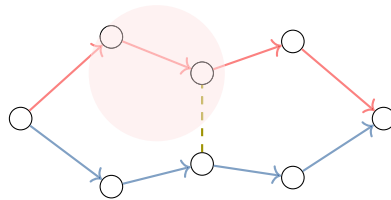
# 1 Introduction

The availability of telecommunications is critical for many infrastructures, people and services, and thus the resilience of telecommunication networks is of utmost importance.

A classical strategy for reaching resilience is to consider routes that are disjoint (either in terms of nodes or links). This is a widely spread technique to provide alternative routing by means of backup paths able of replacing active paths subject to disruptions [6, 11]. The problem of finding  $K$  shortest disjoint paths between two nodes, for a given integer  $K$ , was studied in [1, 10]. When disjoint routes do not exist, solutions are often found by minimizing the extension of the common parts of the routes. This problem is known as the dissimilar and shortest paths problem [8].

However, imposing disjointness, let alone dissimilarity, may not be enough, when a whole area is affected by a disaster. Natural disasters may impact on large areas of a terrain and thus compromise telecommunications routing. Such situation is illustrated in Fig. 1, where the area shaded in red represents the possible impact of a disaster, in this case affecting only part of the top (red) path. In this case the routing could still be recovered by using the bottom (blue) path instead. Due to its relevance, methods to promote the resilience of telecommunication networks against such type of events have been investigated. A recent review on related topics is provided in [9].

Geodiversification is a path protection strategy aimed at ensuring the routing in telecommunication networks in the case of a large scale natural disaster. It is implemented by considering alternative paths for routing if such a disruption affects some nodes, or arcs, of the active path. The availability of an alternative path is ensured by imposing a minimum geographical distance between the paths, which are then said to be geodiverse. The problem of finding geodiverse shortest paths with minimum cost was addressed in [2, 3, 5]. The first of these works addresses the  $K$  geodiverse shortest paths problem by means of a two-steps algorithm (TSA) and two iterative heuristics (iWPSP and MLW). The first step of the TSA algorithm generates sets of shortest disjoint paths, and the second step selects a set of those paths that is sufficiently geodiverse. As for the iWPSP and the MWL heuristics, both consist in the sequential determination of paths that deviate from the shortest by a given minimum distance. In the first case, the deviation is calculated by imposing a central node that is sufficiently distant from the current path and then calculating



**Fig. 1** Geographically diverse paths

the shortest paths from  $s$  to that node and from that node to  $t$ . In the second, the geodiversity is imposed by penalizing the costs of the arcs that are too close to the current path. The heuristics are tested for the case  $K = 2$ . In [3] the authors extend the previous work by considering a delay-skew requirement when using geodiverse paths under area-based challenges. In addition, a flow-diverse minimum-cost routing problem is formulated as a multicommodity flow problem and the trade-off between the delay and skew in choosing geodiverse paths is investigated. In [4, 5], integer linear programming formulations are proposed for finding pairs of geodiverse paths. Both models ensure geodiversity by acting on all pairs of links whose distance from one another is less than  $D$ , and limiting their presence in the solution to one. In [5] it is also addressed the problem of determining the maximum distance between any two paths between the same pair of nodes in a network. A similar problem was also studied in [7] based on a shielding approach.

In the present work, we extend the application of geodiversification found in past literature for  $K = 2$ , and find  $K > 2$  paths between a source and a destination nodes, which are able to ensure end-to-end routing in case of a disaster that affects a circular geographic area with a given radius  $D$ . We address the problem by presenting an integer linear programming (ILP) formulation which allows to obtain  $K$  paths with minimum total cost and that are  $D$ -geodiverse.

In the next section we introduce elementary concepts, formalize the problem and present the ILP formulation. In Sects. 3 and 4, we report the results of preliminary computational experiments obtained by the introduced ILP formulation for instances based on reference networks in telecommunications, the Germany50 and the CORONET CONUS networks, and conclude the work with some final remarks.

## 2 Modeling Approach

Let  $G = (N, A)$  be a geographic network with  $|N| = n$  nodes corresponding to points/locations,  $|A| = m$  arcs, and let  $c_{ij} > 0$  represent the cost associated with the arc  $(i, j) \in A$ . Consider also given initial and terminal nodes,  $s, t \in N$ . The distance between two arcs  $(i, j), (u, v)$  in this network is defined as

$$d_{ij}^{uv} = \min\{d_a^b : a \text{ is a point along } (i, j) \wedge b \text{ is a point along } (u, v)\}$$

where  $d_a^b$  stands for the distance between the points  $a, b$ . That is, this is the smallest distance between the points of the two arcs. Then, the distance between any two paths,  $p, q$ , from  $s$  to  $t$  in this network is defined as

$$d(p, q) = \min\{d_{ij}^{uv} : (i, j) \in p \wedge (u, v) \in q \wedge i, u \neq s \wedge j, v \neq t\}.$$

Alternative definitions of distance between arcs and of distance between paths have been used in the literature (see, for example, [5]). Nevertheless, one argument

for considering only the intermediate components of the paths is that it is not possible to guarantee a minimum distance between the arcs incident on the source or on the destination node. Moreover, it is reasonable to assume that the nodes  $s, t$  are subject to higher protection than the remaining ones, because if a disaster impacts one of them, all communications would be compromised regardless of the number of paths linking them.

Given  $D > 0$ , the paths  $p$  and  $q$  are said to be  $D$ -geodiverse whenever the distance between them is at least  $D$ , that is,

$$d(p, q) \geq D.$$

The dashed line in Fig. 1 illustrates the distance between the two paths linking node  $s$  and  $t$ . The concept can be extended for more than two paths. Let  $K \geq 2$ , a set of  $K$  paths is said to be  $D$ -geodiverse if all pairs of paths within the set are  $D$ -geodiverse.

Let  $X$  denote the set of all sets with  $K$  paths from node  $s$  to node  $t$  in  $G$ . The goal of the  $K$   $D$ -geodiverse shortest paths problem (hereafter KD-GPATHS) is to find  $x \in X$  of minimum total cost, such that the distance between all pairs of paths in  $x$  is at least  $D > 0$ .

Consider the topological binary variables  $x_{ij}^k$ , such that  $x_{ij}^k = 1$  if  $(i, j) \in A$  is in the  $k$ -path, or  $x_{ij}^k = 0$  otherwise, for  $i \in N, k = 1, \dots, K$ . Using these variables, the set  $X$  can be modelled by constraints:

$$\sum_{j \in \delta^+(i)} x_{ij}^k - \sum_{j \in \delta^-(i)} x_{ji}^k = \begin{cases} 1, & i = s \\ 0, & i \in N \setminus \{s, t\} \\ -1, & i = t \end{cases}, \quad k = 1, \dots, K \quad (1)$$

$$x_{ij}^k \in \{0, 1\}, \quad (i, j) \in A, k = 1, \dots, K$$

where  $\delta^+(i), \delta^-(i)$ , is the set of arcs outgoing from, emerging in, node  $i \in N$ , respectively. This set of constraints are flow conservation constraints and ensure that a path from  $s$  to  $t$  exists, for each  $k = 1, \dots, K$ .

Given an arc  $(i, j) \in A$ , let  $N_{ij}^D$  be the set of arcs whose distance from  $(i, j)$  is less than  $D$ . The set  $N_{ij}^D$  can be defined as follows:

$$N_{ij}^D = \{(u, v) \in A' : d_{ij}^{uv} < D\}, \quad (2)$$

where  $A' = A \setminus (\delta^+(s) \cup \delta^-(t))$ .

Then,  $N_{ij}^D$  identifies the arcs that cannot be used if  $(i, j)$  is in the solution. Consequently, using again the previous variables together with this definition, the incompatibility constraints can be modelled as:

$$\sum_{(u,v) \in N_{ij}^D} \sum_{\substack{k'=1 \\ k' \neq k}}^K x_{uv}^{k'} \leq M(1 - x_{ij}^k), \quad (i, j) \in A, \quad k = 1, \dots, K \quad (3)$$

where the constant  $M$  should satisfy

$$M \geq \max_{(i,j) \in A} \{|N_{ij}^D|\}.$$

This prevents the distance between any two paths in  $x$  from being smaller than  $D$ . Thus, the KD-GPaths problem can be modelled by the formulation:

$$\begin{aligned} \min \quad & z = \sum_{(i,j) \in A} c_{ij} \sum_{k=1}^K x_{ij}^k \\ \text{such that} \quad & (1), (3). \end{aligned} \tag{4}$$

### 3 Computational Experiments

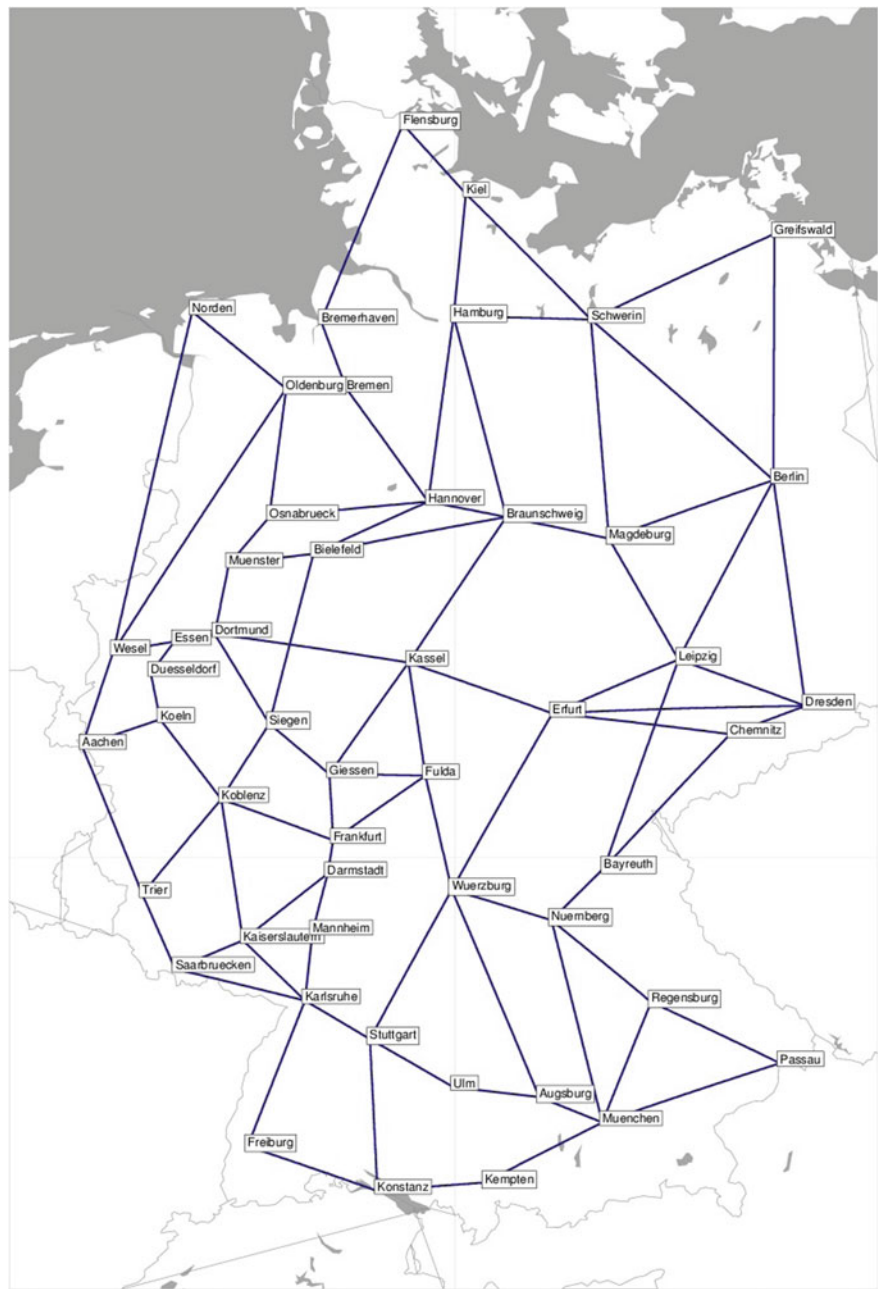
The model presented in the previous section was analyzed empirically with the goal of characterizing the solutions it produces and the associated running times. The experiments were carried out on a 64-bit PC with an Intel®Core™ i9-10900K Quad core with a 3.7GHz processor, with 128GB of RAM. The IBM ILOG CPLEX version 20.1 was the ILP solver used, with no time limit.

The topologies used were the Germany50 and the CORONET CONUS networks (information available at [www.av.it.pt/asou/geodiverse.htm](http://www.av.it.pt/asou/geodiverse.htm)), shown in Figs. 2 and 3, respectively. The Germany50 network has 50 nodes connected by 88 links and the CORONET CONUS network has 75 nodes connected by 99 links. Both networks are undirected, therefore directed networks had to be created based on these ones before solving the formulation. For this purpose, each edge  $\{i, j\}$  was duplicated as the arcs  $(i, j)$  and  $(j, i)$ . The average node degree of the directed versions of the Germany50 and the CORONET CONUS networks are 3.52 and 2.64, respectively. In both cases the degree of the nodes varies between 2 and 5.

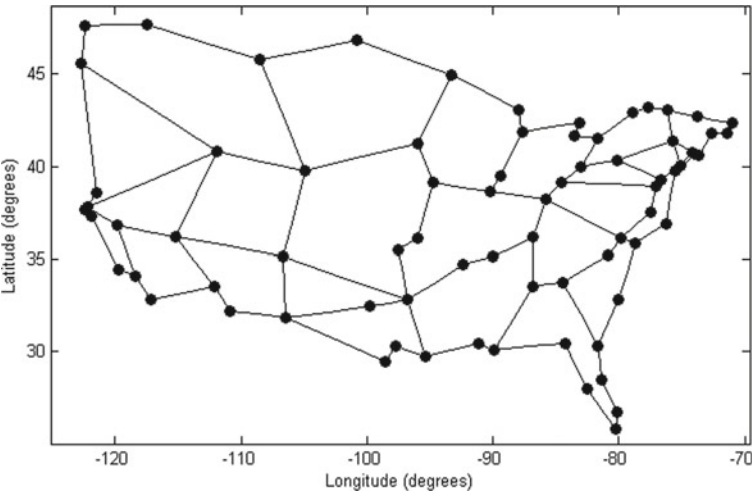
For any arc  $(i, j) \in A$ , the used cost value  $c_{ij}$  was the length of this arc considering the shortest path over the terrestrial surface.

The topological and geographical data, including the distances between the arcs in the networks, was obtained from the repository [www.av.it.pt/asou/geodiverse.htm](http://www.av.it.pt/asou/geodiverse.htm).

In the computational tests, we considered the calculation of  $K = 2, 3, 4, 5$  paths for all source-destination node pairs in the networks and assumed the minimum geographical distances of  $D = 40, 80$  km, for the Germany50 network and  $D = 100, 200$  km for the CORONET CONUS network, given that the latter covers a wider territory than the first. In this way, a total of 9800 instances related to the Germany50 network (1225 per pair  $s - t$ ), and 22 200 instances related to the CORONET CONUS network (2775 per pair  $s - t$ ), were undertaken. The results obtained by (4) are summarized in Table 1. As mentioned earlier, the paths involving only arcs incident in  $s$  or in  $t$  are not subject to the geodiversity constraints. Those cases were omitted from the analysis of the results, in order to allow sounder conclusions.



**Fig. 2** Germany50 network (Retrieved from sndlib.zib.de on February 15, 2023)



**Fig. 3** CORONET CONUS network (Retrieved from <http://www.av.it.pt/asou/geodiverse.htm> on May 19, 2023)

**Table 1** Results obtained by (4)

|                    |               | $D \setminus K$ | 2       | 3       | 4         | 5       |
|--------------------|---------------|-----------------|---------|---------|-----------|---------|
| $N$                | Germany50     | 40              | 1100    | 607     | 142       | 7       |
|                    |               | 80              | 1017    | 302     | 26        | 0       |
|                    | CORONET CONUS | 100             | 2526    | 580     | 24        | 0       |
|                    |               | 200             | 2075    | 328     | 12        | 0       |
| $\bar{z}$          | Germany50     | 40              | 956.51  | 1618.12 | 2322.94   | 3153.00 |
|                    |               | 80              | 1018.46 | 1675.57 | 2286.88   | –       |
|                    | CORONET CONUS | 100             | 5435.28 | 9607.75 | 15 199.54 | –       |
|                    |               | 200             | 5458.54 | 9454.47 | 14 089.17 | –       |
| $\overline{T}$ (s) | Germany50     | 40              | 0.04    | 0.02    | 0.07      | 0.09    |
|                    |               | 80              | 0.04    | 0.07    | 0.10      | –       |
|                    | CORONET CONUS | 100             | 0.02    | 0.02    | 0.04      | –       |
|                    |               | 200             | 0.03    | 0.04    | 0.05      | –       |

$N$ : Number of problems;  $\bar{z}$ : average optimal objective function value;  $\overline{T}$ : average run time (in seconds)

Not surprisingly, the number of solved problems reported in Table 1 decreases with the values of  $K$  and  $D$ . Regarding these results, two features should be highlighted: i) the impact on the results of the increase in the value of  $K$  is much stronger than the impact on the results of the increase in the value of  $D$ , and ii) the decrease in the number of solved problems is especially sharp for  $K > 3$ , in particular for the CORONET CONUS. Naturally, these remarks are instance dependent. The average degree of both considered topologies is small and Figs. 2 and 3 show that nodes with degree greater than 3 are rare.

**Table 2** Average optimal objective function value for the problems that are feasible for all values of  $K$ 

|               | $D \setminus K$ | 2       | 3       | 4        | $N$ |
|---------------|-----------------|---------|---------|----------|-----|
| Germany50     | 40              | 612.11  | 1105.74 | 1976.63  | 26  |
|               | 80              | 643.12  | 1198.85 | 2286.89  | 26  |
| CORONET CONUS | 100             | 3750.62 | 7212.38 | 13005.69 | 12  |
|               | 200             | 3750.62 | 7212.38 | 13005.69 | 12  |

$N$ : Number of problems

On the contrary, the values of the objective function vary in a less expected way, as they decrease as  $D$  rises for the more demanding problems. For each network and each pair  $(K, D)$ , the reported  $\bar{z}$  value corresponds to the average optimal cost for the set of problems solved for that pair. Each of those sets is associated to a different set of pairs  $s - t$ , which hampers the analysis of the results. To overcome this issue, we also study the results for the set of instances that, for each network, were feasible for all the pairs  $(K, D)$ .

For the Germany50 network, the value of  $\bar{z}$  increases with both  $K$  and  $D$ , as expected. However, for the CORONET CONUS network, the value of  $\bar{z}$  does not change when  $D$  increases. This indicates that the  $K$  shortest paths obtained for  $D = 100$  are also feasible (and therefore optimal) for  $D = 200$ . It can also be observed that the values reported in Table 2 are significantly smaller than the ones in Table 1. Thus, we are led to the conclusion that the remaining paths are the shortest ones. Figures 2 and 3 show that the locations with the highest number of connections are close to each other which supports the latter observation.

The determination of the biggest value  $D$  for which the KD-GPaths is feasible would bring complementary information on the impact of the value of  $D$  on the results [5]. Finally, it is worth mentioning that the computational times were always inferior to 0.10 s, therefore the proposed method is efficient and scalable.

## 4 Concluding Remarks

This work addressed the  $K$   $D$ -geodiverse shortest paths problem, with  $K \geq 2$ . An ILP model of the problem was introduced and its performance was studied through a set of computational experiments performed over two real networks taken from the literature (the Germany50 and CORONET CONUS networks) for  $K = 2, \dots, 5$ ,  $D = 40, 80$  km and  $D = 100, 200$  km, respectively, and all possible origin-destination pairs, originating 9800 and 22200 different problems, respectively. Results regarding the percentage of non-feasible instances, the average of the objective function values and required run times were discussed. Two main conclusions can be highlighted:

- the new model was able to solve the problem efficiently for two networks studied;
- 4 or more geodiverse paths between two nodes are very rare in both the topological Germany50 and CORONET CONUS networks.

Future work on this topic should include extending the computational experiments to other topological networks, for instance using other instances in the database [sndlib.zib.de](http://sndlib.zib.de). Also, for a better understanding of the impact of the values of  $D$  on the percentage of non-feasible instances, the problem of the determination of the biggest value  $D$  for which a solution of the KD-GPATHS problem still exists should be investigated, for each  $K$ . Another extension of interest is to study the impact of considering alternative metrics to the distance between paths. For instance, a metric which is able to distinguish the distances between pairs of paths with no intermediate arcs, that is, when one or both paths are composed by at most two arcs.

**Acknowledgements** This work was partially financially supported by the Portuguese Foundation for Science and Technology (FCT) under project grants UIDB/04561/2020, UID/MAT/00324/2020 and UID/MULTI/00308/2020.

## References

1. Bhandari, R.: Survivable Networks: Algorithms for Diverse Routing. Kluwer Academic Publishers, USA (1998)
2. Cheng, Y., Gardner, M.T., Li, J., May, R., Medhi, D., Sterbenz, J.: Analysing geopath diversity and improving routing performance in optical networks. *Comput. Netw.* **82**, 50–67 (2015)
3. Cheng, Y., Medhi, D., Sterbenz, J.: Geodiverse routing with path delay and skew requirement under area-based challenges. *Networks* **66**, 335–346 (2015)
4. de Sousa, A., Santos, D.: The minimum cost  $D$ -geodiverse anycast routing with optimal selection of anycast nodes. In: 2019 15th International Conference on the Design of Reliable Communication Networks (DRCN), pp. 21–28 (2019)
5. de Sousa, A., Santos, D., Monteiro, P.: Determination of the minimum cost pair of  $D$ -geodiverse paths. In: DRCN 2017-Design of Reliable Communication Networks; 13th International Conference, pp. 1–8 (2017)
6. Gomes, T., Craveirinha, J., Jorge, L.: An effective algorithm for obtaining the minimal cost pair of disjoint paths with dual arc costs. *Comput. Oper. Res.* **36**, 1670–1682 (2009)

7. Liu, M., Qi, X., Pan, H.: Optimizing communication network geodiversity for disaster resilience through shielding approach. *Reliab. Eng. Syst. Saf.* **228**, 108800 (2022)
8. Moghanni, A., Pascoal, M., Godinho, M.T.: Finding  $K$  shortest and dissimilar paths. *Int. Trans. Oper. Res.* **29**, 1573–1601 (2022)
9. Rak, J., Girão-Silva, R., Gomes, T., Ellinas, G., Kantarci, B., Tornatore, M.: Disaster resilience of optical networks: state of the art, challenges, and opportunities. *Opt. Switch. Netw.* **42**, 100619 (2021)
10. Suurballe, J.: Disjoint paths in a network. *Networks* **4**, 125–145 (1974)
11. Xu, D., Chen, Y., Xiong, Y., Qiao, C., He, X.: On finding disjoint paths in single and dual link cost networks. In *IEEE INFOCOM 1*, 2004 (2004)

# Are the Portuguese Fire Departments Well Located for Fighting Urban Fires?



Maria Isabel Gomes<sup>ID</sup>, Ana C. Jóia, João Pinela,  
and Nelson Chibeles-Martins<sup>ID</sup>

**Abstract** In Portugal, there has been a lack of strategic analysis regarding the preparedness level of urban firefighting in the current Fire Department (FD) locations. The opening of several FDs was the result of the desire of the population and the need to have a quick response in case of an incident. Thus, Portugal has a high number of FDs in relation to the small size of the country and an uneven distribution. Therefore, it is impossible to guarantee an efficient intervention and a timely response throughout the country. In this work, the current situation of the FD locations in the district of Porto is studied. An optimisation approach based on coverage models is carried out. The results show where the FDs' areas of operation overlap, where there is a deficit in coverage and which existing locations are strategically important to maintain for urban fire fighting. Although only the district of Porto was studied, the proposed approach can be applied to all other Portuguese districts.

**Keywords** Location modelling · Coverage optimisation models · Case study · Urban firefighting

## 1 Introduction

The first few minutes after a fire has started in a building are critical to ensuring the safety of people and the protection of property [1]. The ability to quickly evacuate occupants, extinguish the fire and limit damage to the structure is highly dependent on the occupants' response to the fire and the fire service's response time to the call [2]. Response time, the time between the call and arrival at the fire scene, is heavily affected by the number and location of fire departments (FDs). Insufficient number or poor location of FDs can significantly increase response times and lead

---

M. I. Gomes (✉) · N. Chibeles-Martins  
NOVA MATH, Campus da Caparica, 2829-516 Caparica, Portugal  
e-mail: [mirg@fct.unl.pt](mailto:mirg@fct.unl.pt)

M. I. Gomes · A. C. Jóia · J. Pinela · N. Chibeles-Martins  
Department of Mathematics, Campus da Caparica, 2829-516 Caparica, Portugal

to serious loss of life and property [3]. A faster response time results in better fire protection and reduces direct financial losses. Therefore, the ability of FDs to protect the population from fire damage is strongly influenced by their geographical location [4]. Appropriate placement of FDs can reduce extensive overlap between areas to be served by different FDs, thereby maximising the use of human and material resources across the geographical area of operation.

In Portugal, the main driving force behind the opening of FDs has been the population's desire for a rapid response to emergencies (fires, road accidents, floods...). This has led to a seemingly uneven distribution of FDs across the country. At present, there are about 450 FDs, most of which are humanitarian associations and their human resources are largely voluntary. Recently, two studies have dealt with the Portuguese reality in terms of the performance efficiency of the FDs [5], and the spatial modelling of fire events in urban and residential areas [6]. Both studies mention the need for further research on the number and location of FDs.

The study of the spatial distribution of FDs using analytical models dates back to the late 60s, early 70s with seminal works by Hogg [7], Toregas et al., [8] and Guild and Rollin [9]. The problem of locating emergency services facilities, such as FDs, has traditionally been modelled as a coverage problem. The most common models in the literature are the Set Covering Problem (SCP) and the Max Covering Problem (MCP). Both approaches assume the existence of a critical coverage time or distance within which events must be serviced to be considered covered. This time or distance is assumed to be known in advance [10]. The SCP aims to select locations to minimise the number of facilities required while still covering all demand nodes [11]. However, in many situations it is not feasible or even possible to cover all demand nodes. For example, when locating FDs, it may be necessary to open a fire station to 'cover' a house in a fairly remote area. The objective of the MCP is then to locate a predetermined number of facilities that will maximise the number of demand points covered within a given time or distance. In most cases, the number of available facilities will not 'cover' all the demand points, and therefore the MCP will find the locations that maximise demand coverage [12]. This approach has been applied in several studies (e.g. the optimisation of fire station locations for Istanbul Metropolitan Municipality [13], the optimal fire station locations in Nanjing, China [14], or a study on fire service coverage in Brisbane, Australia [15]). For more interested readers, we recommend three reviews of location models for emergency services: [3, 16, 17].

The above models deal with primary location issues and therefore assume that the data is deterministic, which may lead to an oversimplification of the problem. To address this drawback, some authors have approached the location decision using discrete location models, address the uncertainty and more tactical nature decision as vehicle and human resource allocation to fire events separately. Chevalier et al. [18], develop a decision support system to help Belgian authorities understand the functioning of their FDs. As emergency events, the authors considered not only urban fires, but also non-urban fires, medical emergencies, non-medical emergencies and non-emergency events. The location problem is modelled with an extension of the Set Covering Problem in order to distinguish events according to their frequency and legal requirements. The hypercube queuing model was applied to solve the staffing

problem (the allocation of human resources to events). Another interesting work was developed by Rodriguez et al. [19], which dealt with the optimal location of fire stations in the Concepción region of Chile, as well as the allocation of vehicles to each fire station. The optimal locations are obtained by maximising the coverage of expected emergencies under a response time constraint, the possibility of relocating existing fire stations or installing new ones. It also considers the possibility of deploying different numbers and types of vehicles depending on the region and the type of emergency. The modelling approach was inspired by the work of Chevalier et al. [18]. Later, the authors refined the work by developing an iterative simulation-optimisation approach, using a robust formulation for the location problem and a discrete event simulation approach for the vehicle allocation decisions [20].

To the best of our knowledge, this is the first optimisation-based study to consider the location of Portuguese FDs in relation to their level of preparedness for urban firefighting. It uses real data from urban fire events provided by the Portuguese National Emergency and Civil Protection Agency (ANEPC). We aim to (i) study which are the optimal locations for FDs if a strategic study had been carried out; (ii) understand, based on the current locations, which areas have a coverage deficit and where new FDs should be opened; and (iii) determine which of the existing FDs are currently indispensable for timely response to urban fire events, as it is suspected that some areas have more FDs than needed. In terms of uncertainty, we have developed a small simulation model to evaluate the locations proposed by the optimisation models. However, this model and the corresponding analysis of results are not included in this paper due to space limitations.

This paper is structured as follows. Section 2 describes the Portuguese firefighting network and urban fire data, providing details that are fundamental to understanding the characteristics of the problem. The modelling approach is presented in Sect. 3. In particular, the two linear optimisation models are presented. The corresponding results are then discussed (Sect. 4). The paper ends with some concluding remarks and directions for future work (Sect. 5).

## 2 Case Study

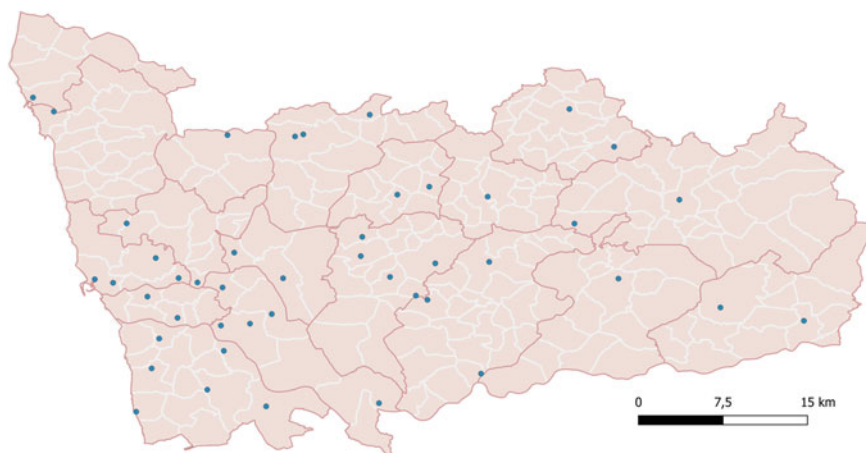
In Portugal, buildings and premises are classified into four different risk categories according to their type of use (dwellings, schools, hospitals...) [21]. This risk classification is based on several factors: (i) the height of the building, (ii) the estimated maximum number of people that can occupy a given space of a building or enclosure at the same time, (iii) the modified fire load density, and (iv) the presence of floors below the reference level, among others [22]. Category I is that of the lowest risk buildings, while Category IV is that of the highest risk buildings. Fire events in buildings of the latter category will have a very low probability of occurrence but a very high impact. The reference values for the level of rescue preparedness depend on factors such as “(a) distance and maximum travel time, by normal access routes, between the fire department and the type of use of the building or enclosure;

(b) technical means (vehicles and equipment) that can be mobilised for immediate deployment after the alarm; and (c) human resources, in minimum numbers (minimum operational intervention force), on duty 24 h a day to operate the technical means mentioned in the previous paragraph...". However, these limits have been established only for buildings and premises in risk categories III and IV: "the maximum distance to be travelled between the facilities of a fire brigade should be up to ten kilometers, this distance being based on compliance with a maximum travel time, at an average speed of 55 km/h, of approximately ten minutes after the dispatch of the first alarm" [23].

The ANEPC database provided the date, time and location (in latitude and longitude coordinates) of each urban fire event. The coordinates available varied in precision, with the smallest unit of precision common to all records being the "freguesia" or civil parish, a subdivision of the municipality (hereafter referred to as "parish"). The total annual number of urban fires in Portuguese municipalities varied between 8,139 and 9,832 in 2014 and 2019 respectively, with an annual average of around 8,840 events [6]. It was not possible to assess the risk category associated with the events and therefore, in agreement with the ANEPC experts, all the events in the database were considered to belong to risk categories I and II—frequent events. Consequently, there is no real data on risk categories III and IV—rare events. In order to model the possibility of fires in the latter risk category, it was assumed that parishes where rare events could occur should be covered by at least one FD, i.e. at least one FD should be located within a 10 min radius. For the frequent events, it was assumed that a fraction of their total number should be covered. Note that this does not mean that some fire events will go unattended. What it does mean is that it may take more than 10 min for the first brigade to reach them.

The municipality is the predefined geographical area in which the FD regularly operates and/or is responsible for the first intervention. According to Portuguese law, the area of operation corresponds to the municipality in which the FD is located (if there is only one FD) or it is a geographical area that must coincide with one or more contiguous parishes (if there is more than one FD operating in the municipality). This administrative organisation rule does not ensure that fire events are assigned to the closest FD, but rather to the one located in the same municipality. In this study, this rule will be relaxed, as we want to understand which are the most interesting locations for FDs that guarantee the level of rescue preparedness set by the legislation. As a consequence, all the analyses carried out on the current location of the FDs will show higher levels of rescue preparedness than the real one. Nevertheless, it provides valuable information for the strategic analysis of the location of Portuguese FDs.

The statistical analysis of the data is beyond the scope of this paper. However, in order to understand why we have chosen the data used throughout the paper, it is important to summarise the main conclusions regarding temporal and spatial dependencies. Bispo et al. [6], applied appropriate analysis techniques and concluded that there is no evidence of temporal dependence between the annual data, but there is a significant spatial dependence across the country (Moran's I and Geary's C global measures of autocorrelation indicate significant spatial clustering with  $p < 0.001$ ). Furthermore, the authors showed that the occurrence of urban fires is clearly city-



**Fig. 1** FD locations in Porto district

centred around the main cities of Lisbon and Porto. In particular, the municipality of Porto is of particular interest, as it has a very high building density and one of the highest fire incidences. The authors also conclude that the metropolitan area of Porto has the highest association with spatial neighbours and that municipalities within the district of Porto have one of the highest predicted incidences of urban fires per  $\text{km}^2$  (fitted spatial *Durbin error* model, explaining around 97% of the variability in the data). Therefore, the district of Porto and the most recent year available in the database (2020) were chosen as the location and time period for our study.

Figure 1 shows the district of Porto, the boundaries of its 243 parishes and the location (with a higher degree of geographical precision than the parish level) of the 45 FDs currently operating in this district. The FDs are spread over 42 parishes, which means that 3 parishes have 2 FDs. Each parish is used as a location for a fire event and as a potential location for an FD. Information on each parish has been summarised at a reference point corresponding to the location of the *Junta de Freguesia* - the administrative building of the parish. Centroids are traditionally the reference points when discretising a continuous area, but for more rural parishes this could mean that information was aggregated at a point located, for example, in the middle of an agricultural field with no buildings surrounding it. We therefore decided to use a more or less central point in the urban area of the parish.

### 3 Modelling Approach

As mentioned above, the two models used to strategically study the optimal locations of the FDs are the SCP and the MCP. The former with slight modifications of the classical formulation. Although both models are well known in the location com-

munity, for the sake of completeness of this work we have chosen to present both formulations.

### *Set Covering Model*

The SCP model is based on the model proposed in [18]. Let  $N$  be the set of all parishes (for both events and FDs possible locations). Since the exact location of existing FDs is known, let  $G$  be the set of existing FDs such that  $N \cap G = \emptyset$ . The set of all locations is defined as  $L = N \cup G$ . Given the need to ensure that the response time to any rare event is met, let  $S \subseteq N$  be the set of all parishes where there is at least one risk category III or IV building. Let  $\bar{T}$  be the maximum response time (or distance) and define  $T_j = \{i \in L : t(i, j) \leq \bar{T}, j \in N\}$  where  $t(i, j)$  is the time (or distance) to travel between locations  $i$  and  $j$  (the neighbourhood set). In terms of data, define  $\alpha$  as the fraction of frequent events to be covered and  $w_j$  as the number of (frequent) events that have occurred in parish  $j$ . Two sets of binary variables are needed:  $X_i = 1$  if an FD is located in parish  $i$ ,  $i \in L$ , and  $= 0$  otherwise; and  $Y_j = 1$  if parish  $j$  is covered by at least one FD, and  $= 0$  otherwise. The model that minimises the number of FDs needed to cover all parishes with high-risk buildings and at least a fraction of the frequent events are covered in at most  $\bar{T}$  units, is given by:

$$\min \sum_{i \in L} X_i - \varepsilon \sum_{i \in N} w_i Y_i \quad (1)$$

$$\text{s.t.} \quad \sum_{i \in T_j} X_i \geq Y_j, \quad j \in N \quad (2)$$

$$\sum_{j \in N} w_j Y_j \geq \alpha \sum_{j \in N} w_j \quad (3)$$

$$\sum_{i \in T_j} X_i \geq 1, \quad j \in S \quad (4)$$

$$\begin{aligned} X_i &\in \{0, 1\} \quad i \in L \\ Y_j &\in \{0, 1\} \quad j \in N \end{aligned} \quad (5)$$

The objective function is defined by the Eq. (1) and minimises the total number of open FD locations (first term). The second term is a penalty factor that “encourages” the coverage of parishes with a higher number of fire events, where  $\varepsilon > 0$  is small enough. Constraint (2) ensures that parish  $j$  is only considered covered if at least one FD can reach it in at most  $\bar{T}$  units (the maximum response time). Constraint (3) sets the minimum number of frequent events that must be covered. Constraint (4) ensures that all parishes with at least one risk category III or IV building are covered. Finally, constraints (5) set variables domain.

This model does not distinguish between the different response times for buildings in higher and lower risk categories. This can be easily modelled by considering two different values for the maximum response time and defining the corresponding neighbourhood sets.

### ***Maximal Covering Location Model***

As proposed by Church and Reville (1974), the Maximal Covering Location model aims to locate  $Q$  facilities covering the maximum amount of demand (fire events in our case) [24].

$$\max \sum_{j \in N} w_j Y_j \quad (6)$$

$$\text{s.a.} \sum_{i \in T_j} X_i \geq Y_j, \quad j \in N \quad (7)$$

$$\sum_{i \in L} X_i = Q \quad (8)$$

$$\begin{aligned} X_i &\in \{0, 1\} \quad i \in L \\ Y_j &\in \{0, 1\} \quad j \in N \end{aligned} \quad (9)$$

The objective function (6) maximizes the number of covered frequent events, i.e., maximizes the expected coverage. The constraints (7) ensure that each parish  $j$  is considered covered if at least one FD from its neighbourhood set is opened. Equation (8) specifies the number of FD to open. Finally, constraints (9) set variables domain.

In this study, the number of facilities resulting from the coverage of the SCP model is used to determine the value of  $Q$ . Therefore, the model searches for the exact  $Q$  FD locations that cover the largest number of parishes. If, for a given value of  $Q$  FDs, the total number of fire events covered within the maximum response time ( $\bar{T}$  units) is less than 100%, this means that fires occurring in those parishes will be rescued with a response time longer than that required by legislation.

## **4 Results and Discussion**

In order to study the optimal locations for FDs in the district of Porto, four scenarios have been established, which differ in the way the existing FD locations have been considered. The SCP model is applied to each scenario and the results obtained are discussed and compared. For some of the scenarios and assuming the number of FDs proposed by the SCP model, the MCL model is solved and the results are discussed and compared. It should be noted that all conclusions drawn below consider urban firefighting as the sole activity of the FDs. They should therefore be read with the understanding that a new analysis should be carried out if all the other activities carried out by the FDs are taken into account.

## 4.1 Scenarios

Scenario 1 examines the location of FDs as if none had been established, the idyllic scenario. This allows us to understand what is the minimum number of FDs and their best location to ensure full coverage of parishes where rare events may occur and a pre-determined fraction of coverage for frequent events. Scenario 2 looks at which new sites should be added to the 45 existing FDs. This scenario will indicate which areas are not currently covered. Scenario 3 takes all the locations from scenario 2 and will investigate which existing FDs are essential to meet the required coverage levels. In other words, it will show which existing FDs are redundant in terms of the level of rescue preparedness for urban firefighting. Finally, in Scenario 4 we consider as possible locations for FDs either the parish reference points or the exact coordinates of the 45 FDs currently in operation. The aim of this scenario is to understand which of the existing sites are essential and which new sites should be added. However, given the possibility of not using an existing site to open a new one nearby (both with similar coverage areas), we give preference to the existing FDs by penalising the objective function whenever a new site is considered. Therefore, a different objective function is considered in scenario 4.

## 4.2 Fraction of Coverage for Frequent Events

Before applying the proposed models to the above scenarios, it is necessary to estimate an appropriate value for the fraction of frequent events that should be covered (parameter  $\alpha$ ). Chevalier et al. assumed a coverage of 90% when carrying out a similar study for Belgium [18]. However, we did not know whether this value would be meaningful in the Portuguese context. Therefore, an annual analysis was carried out, taking into account the current location of the FDs. Table 1 shows the maximum percentage of frequent events that can be covered by the number and location of the current FDs.

The maximum percentage of frequent events covered varies between 89% and 92%. It should be noted that these values have been determined taking into account the best possible coverage, which may not correspond to the real coverage. Regarding the coverage of locations where rare events may occur, 11 out of 73 parishes where this type of event may occur are not covered by the 45 FDs currently operating in the district of Porto. Therefore, although the existing FDs allow a coverage of about 90% of frequent events, their locations do not ensure the legal requirements for rare

**Table 1** Fraction of coverage for frequent events, per year

| 2012  | 2013  | 2014  | 2015  | 2016  | 2017  | 2018  | 2019  | 2020  |
|-------|-------|-------|-------|-------|-------|-------|-------|-------|
| 91.12 | 90.21 | 90.26 | 91.31 | 89.85 | 90.17 | 91.08 | 91.98 | 91.31 |

events imposed by the legislation. Therefore, the four scenarios presented below will allow the study of the locations of the FDs so that, while ensuring that at least 90% of frequent events are covered, all rare events are covered together.

### 4.3 *Best Locations and Optimal Number of FDs*

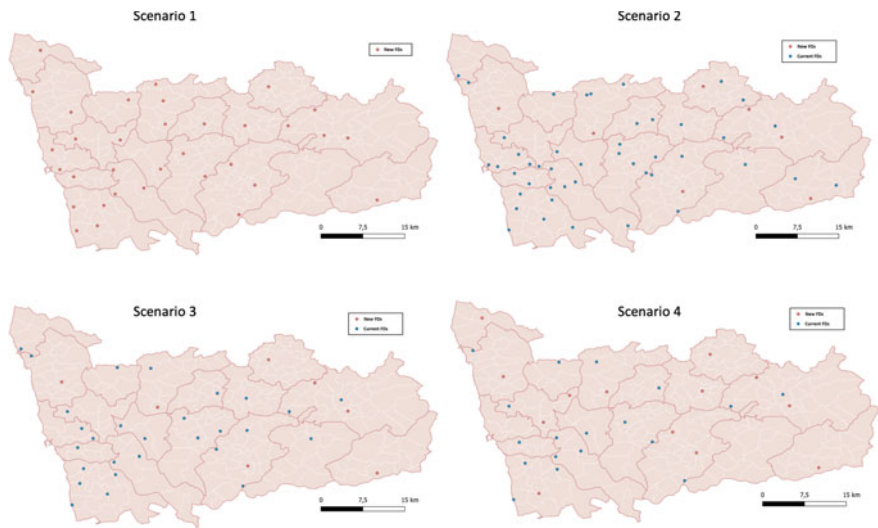
The optimal number of FDs varies between 31 and 52 (see Table 2). For scenario 1, the opening of 33 FDs is proposed. Their locations are distributed throughout the district, being more concentrated in parishes with a higher population density, those closer to the city of Porto (Fig. 2, scenario 1). Taking into account the existing 45 FDs, seven more would have to be installed to cover the parishes where rare events could occur. Figure 2, scenario 2, shows the current locations (blue dots) and where new FDs should be placed (red dots). The later locations are placed in areas where there are uncovered parishes. If only the locations obtained in scenario 2 are considered as possible locations (scenario 3), then only 28 of the 45 existing FDs need to be maintained. With the seven new FDs, a total of 35 FDs will provide adequate coverage (Fig. 2, scenario 3). Finally, taking into account both the reference points and the locations of the existing FDs, only 31 FDs are required: 18 of the existing ones and 13 new ones (Fig. 2, scenario 4). The fact that the two models have different objective functions explains the lower number of FDs than that proposed in scenario 1.

### 4.4 *Sensitivity Analysis of the Coverage Fraction for Frequent Events*

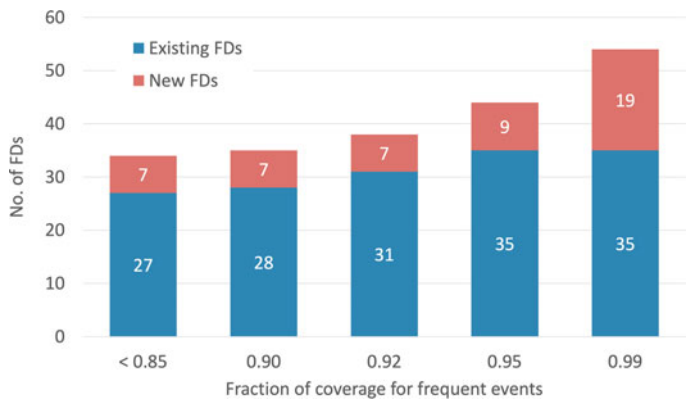
In order to understand the impact of assuming a coverage rate of 90%, a sensitivity analysis was carried out in the context of scenarios 2 and 3. These are the scenarios that best capture a possible strategy for the reorganisation of the Portuguese FD network. If this proportion is less than 93%, the number and location of FDs will not change for scenario 2. The need for total coverage of the rare events provides a coverage level of 93% for the frequent events. For scenario 3 (Fig. 3), changes to this parameter result in a different network configuration due to the more stringent

**Table 2** Minimum number of FDs needed for each scenario

|              | Scenario 1 | Scenario 2 | Scenario 3 | Scenario 4 |
|--------------|------------|------------|------------|------------|
| New FDs      | 33         | 7          | 7          | 13         |
| Existing FDs | –          | 45         | 28         | 18         |
| Total        | 33         | 52         | 35         | 31         |

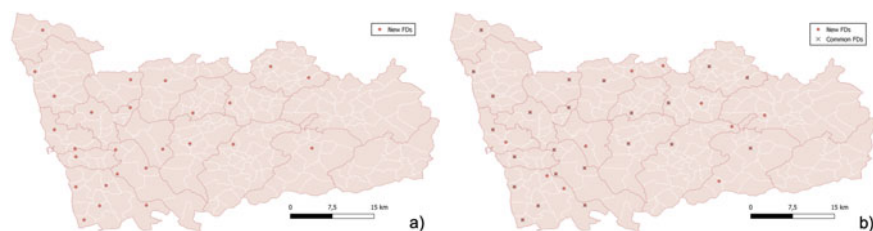


**Fig. 2** Locations of FDs for each scenario



**Fig. 3** Number of FDs for different levels of coverage of frequent events in scenario 3

number of FDs available. However, the same 7 new FDs are required when this level reaches 92%. Below this value, changes will only affect the existing FDs. The maximum number of existing FDs required in this analysis is 35 (10 less than the current FDs). Above 95%, the lack of coverage is overcome by opening new FDs. Note that to achieve these levels of coverage for frequent events, FDs are opened in areas with very low population density. Of the 28 existing FDs that are maintained with a coverage rate of 0.90, 26 locations are common to all values tested.



**Fig. 4** FDs optimal locations for frequent event coverage **a** SCL model, **b** MCL model

## 4.5 Analysis of Frequent Fires Coverage

For the Maximal Covering Location (MCL) model, it is necessary to set a number of FDs to be located, parameter  $Q$ . This value is the one proposed by the SCL model for the corresponding scenario. Note that rare events are not considered in the MCL model, i.e. it is not guaranteed that all parishes where rare events may occur are covered. There are two reasons for this. Firstly, there are no rare events reported in the data provided by ANEPC. Secondly, given the requirement to have an FD in parishes where there are Category III and IV buildings, since parishes are the smallest unit of location, a very small number of FD locations are “available for relocation” in scenarios 2 and 3. Therefore, for the sake of comparability, the SCL has been solved considering only the frequent events.

Figure 4 shows, side by side, the optimal locations for FDs when only frequent fire events are considered. On the left, Fig. 4a), are the locations suggested by the SCP model. Only 26 FDs are needed to cover all frequent events, which is 7 FDs less than if all parishes with category III and IV buildings were to be covered. 15 of the 26 locations are the same as those proposed by the SCP model when considering the rare events. When setting  $Q = 26$ , the MCL model relocates two FDs (see Fig. 4b). In scenario 2, the 45 current locations provide 91% coverage of frequent fires. Therefore, if rare events are not considered, no additional location is needed, which rendering the MCL model solution meaningless. Finally, if the existing FDs are allowed to “close” (scenario 3), only 29 of the 45 provide 90% coverage of frequent events. In other words, 16 of the current FDs only contribute to the coverage of about 1% of frequent events. Regarding the location of these 29 FDs, 24 of them are in the same location as the one proposed by the SCP when modelling frequent and rare events. This confirms the strategic nature of these locations in terms of the level of preparedness required to fight urban fires.

## 5 Final Remarks

In Portugal, the location of fire stations has not been the subject of a strategic study. The work carried out was aimed at optimising the location of fire stations in the district of Porto, taking into account only events related to urban firefighting. Discrete location models based on coverage are used to minimise the number of FDs in order to guarantee full coverage of rare events and partial coverage of frequent events, in accordance with the current Portuguese legislation. Once the minimum number of FDs required had been determined, the maximum coverage that could be achieved with the same number of fire stations was studied. The results show that there are locations common to the solutions proposed in all the scenarios, making them very important locations in terms of the level of preparedness required to fight urban fires. It has also shown which regions have overlapping FD areas of operation and which have a deficit in coverage.

The proposed approach can easily be applied to the remaining districts of mainland Portugal. Furthermore, although outside the scope of this study, the proposed approach can be used as a tool to evaluate the impact of different funding scenarios, which can be adequately modelled by systematically changing the number of FDs and solving the corresponding optimisation model.

While the main driving force behind the establishment of Fire Departments in Portugal has been the population's desire for a rapid response to emergencies, it is important to note that this particular study focuses specifically on preparedness for urban firefighting, as this was the only data we were able to assess. We recognise that other emergencies, such as road traffic accidents and floods, also require an effective response. The methodology used is suitable for studying other types of emergency. To do this, data on these events must be collected and analysed, as was done for this study. Therefore, further research and consideration is needed to expand the scope of activities performed by FDs. It should be noted that if these other events are included, the conclusions drawn regarding the locations of FDs will certainly suffer some, if not, many changes.

Finally, this study did not address uncertainties in the location of these events or in the travel times of the first intervention vehicles, which affect the response time. This is the study that follows the one published in this manuscript.

**Acknowledgements** This work is funded by national funds through the Fundação para a Ciência e a Tecnologia, I.P., under the scope of the projects DSAIPA/DS/0088/2919, and UIDB/00297/2020 and UIDP/00297/2020 (Center for Mathematics and Applications). We would like to acknowledge the invaluable collaboration of Eng. José Pedro Lopes, from ANEPC. Without his support, this work would not have been possible.

## References

- Manes, M., Rush, D.: A comprehensive investigation of the impacts of discovery time and fire brigade response time on life safety and property protection in England. *Indoor Built Environ.* **32**(1), 242–261 (2023)
- Swan, D., Baumstark, L.: Does every minute really count? Road time as an indicator for the economic value of emergency medical services. *Value in Health* **25**(3), 400–408 (2022)
- Wang, W., Wu, S., Wang, S., Zhen, L., Qu, X.: Emergency facility location problems in logistics: status and perspectives. *Transp. Res. part E: Logist. Transp. Rev.* **154**, 102465 (2021)
- Yang, L., Jones, B.F., Yang, S.H.: A fuzzy multi-objective programming for optimization of fire station locations through genetic algorithms. *Eur. J. Oper. Res.* **181**(2), 903–915 (2007)
- Eslamzadeh, M.K., Grilo, A., Espadinha-Cruz, P.: A model for fire departments' performance assessment in Portugal. *Fire* **6**(1), 31 (2023)
- Bispo, R., Vieira, F. G., Bachir, N., Espadinha-Cruz, P., Lopes, J.P., Penha, A., Marques, F.J., Grilo, A.: Spatial modelling and mapping of urban fire occurrence in Portugal. *Fire Saf. J.* **103802** (2023)
- Hogg, J.M.: The siting of fire stations. *J. Oper. Res. Soc.* **19**(3), 275–287 (1968)
- Toregas, C., Swain, R., ReVelle, C., Bergman, L.: The location of emergency service facilities. *Oper. Res.* **19**(6), 1363–1373 (1971)
- Guild, R.D., Rollin, J.E.: A fire station placement model. *Fire Technol.* **8**(1), 33–43 (1972)
- Boonmee, C., Arimura, M., Asada, T.: Facility location optimization model for emergency humanitarian logistics. *Int. J. Disaster Risk Reduct.* **24**, 485–498 (2017)
- Daskin, M.S.: *Network and Discrete Location: Models, Algorithms, and Applications*. Wiley, New York (1995)
- Arabani, A.B., Farahani, R.Z.: Facility location dynamics: an overview of classifications and applications. *Comput. Ind. Eng.* **62**(1), 408–420 (2012)
- Aktaş, E., Özyayın, Ö., Bozkaya, B., Ülengin, F.: Önsel, Ş: Optimizing fire station locations for the Istanbul metropolitan municipality. *Interfaces* **43**(3), 240–255 (2013)
- Yao, J., Zhang, X., Murray, A.T.: Location optimization of urban fire stations: access and service coverage. *Comput. Environ. Urban Syst.* **73**, 184–190 (2019)
- Kc, K., Corcoran, J., Chhetri, P.: Spatial optimisation of fire service coverage: a case study of Brisbane. *Aust. Geograph. Res.* **56**(3), 270–284 (2018)
- Simpson, N.C., Hancock, P.G.: Fifty years of operational research and emergency response. *J. Oper. Res. Soc.* **60**, S126–S139 (2017)
- Bólinger, V., Ruiz, A., Soriano, P.: Recent optimization models and trends in location, relocation, and dispatching of emergency medical vehicles. *Eur. J. Oper. Res.* **272**(1), 1–23 (2019)
- Chevalier, P., Thomas, I., Geraets, D., Goetghebeur, E., Janssens, O., Peeters, D., Plastria, F.: Locating fire stations: an integrated approach for Belgium. *Socioecon. Plann. Sci.* **46**(2), 173–182 (2012)
- Rodriguez, S.A., Rodrigo, A., Aguayo, M.M.: A facility location and equipment emplacement technique model with expected coverage for the location of fire stations in the Concepción province. Chile. *Comput. Ind. Eng.* **147**, 106522 (2020)
- Rodriguez, S.A., Rodrigo, A., Aguayo, M.M.: A simulation-optimization approach for the facility location and vehicle assignment problem for firefighters using a loosely coupled spatio-temporal arrival process. *Comput. Ind. Eng.* **157**, 107242 (2021)
- Autoridade Nacional de Proteção Civil: *Segurança contra incêndio em edifícios nota técnica n.º 01 utilizações-tipo de edifícios e recintos*, Portugal (2013)
- Assembleia da República: *Lei n.º 123/2019 de 18 de Outubro*, Portugal (2019)
- Autoridade Nacional de Emergência e Proteção Civil: *Despacho n.º 8903/2020*, Ministério da Administração Interna, Portugal (2020)
- Church, R., ReVelle, C.: The maximal covering location problem. *Pap. Region. Sci. Assoc.* **32**(1), 101–118 (1974)

# Parcel Delivery Services: A Sectorization Approach with Simulation



Cristina Lopes, Ana Maria Rodrigues, Elif Ozturk, José Soeiro Ferreira,  
Ana Catarina Nunes, Pedro Rocha, and Cristina Teles Oliveira

**Abstract** Sectorization problems, also known as districting or territory design, deal with grouping a set of previously defined basic units, such as points or small geographical areas, into a fixed number of sectors or responsibility areas. Usually, there are multiple criteria to be satisfied regarding the geographic characteristics of the territory or the planning purposes. This work addresses a case study of parcel delivery services in the region of Porto, Portugal. Using knowledge about the daily demand in each basic unit (7-digit postal code), the authors analysed data and used it to simulate dynamically new daily demands according to the relative frequency of service in each basic unit and the statistical distribution of the number of parcels to be delivered in each basic unit. The sectorization of the postal codes is solved independently considering two objectives (equilibrium and compactness) using Non-dominated Sorting Genetic Algorithm-II (NSGA-II) implemented in Python.

**Keywords** Sectorization · Delivery services · Simulation · Multiple criteria · Genetic algorithms

---

C. Lopes (✉) · A. M. Rodrigues · C. T. Oliveira  
CEOS.PP, ISCAP, Polytechnic of Porto, Porto, Portugal  
e-mail: [cristinalopes@iscap.ipp.pt](mailto:cristinalopes@iscap.ipp.pt)

A. M. Rodrigues · E. Ozturk · J. S. Ferreira · P. Rocha  
INESC TEC - Technology and Science, Porto, Portugal

E. Ozturk  
FEP.UP - Faculty of Economics, University of Porto, Porto, Portugal

J. S. Ferreira  
FEUP - Faculty of Engineering, University of Porto, Porto, Portugal

A. C. Nunes  
ISCTE - University Institute of Lisbon, Lisbon, Portugal

CMAFcIO - Faculty of Sciences, University of Lisbon, Lisbon, Portugal

## 1 Introduction

Generally, sectorization consists of partitioning a geographic region composed of indivisible basic units into smaller areas (sectors). This division is done according to several criteria, among which the most common are equilibrium, compactness, and contiguity. This means that usually, it is desirable that each sector represents equal parts of the whole, that it has a compact shape, and that it is possible to connect any basic units within the same sector. Sectorization problems appear in a wide range of real problems. In particular, they are used in commercial territory design in order to divide a large geographical area into smaller regions of responsibility. Each salesperson or team is in charge of one small area or sector. This paper introduces a similar situation in delivery services. Data with the daily demand in each 7-digit postal code is known, and the purpose is to define groups of nearby postal codes with balanced demand to assign to each delivery truck. It is assumed that a good knowledge of the locations and routes within the designated postal codes area increases the productivity of the delivery teams. This assumption means that the responsibility areas of each team should be robust. However, as the number of parcels to be delivered in each location varies daily, the sectors should be able to adapt to change. Therefore, this case study is an application of sectorization where the instances have dynamic contours.

The contributions of this paper are addressing a real case study of parcel delivery services in the region of Porto, Portugal. The purpose is to deal with the situation, following a sectorization approach and recurring to Monte Carlo simulation. This approach facilitates the definition of routes afterwards, by dividing a large problem into several smaller problems, ensuring a balanced division of labour among employees.

Solution tests resort to simulation from actual historic data provided by the company. Section 2 presents sectorization, namely within the context of Commercial Territory Design, also referring to an extension to dynamic aspects. Section 3 describes the Parcel Delivery Case Study and how simulation handles larger data and prepares for new difficulties. How the sectorization results are obtained using Non-dominated Sorting Genetic Algorithm-II (NSGA-II) multiple criteria algorithm is the content of Sect. 4. Section 5 concludes the paper and refers to future work.

## 2 Sectorization

Sectorization problems, also known as districting or territory design, appear when trying to group a set of previously defined basic units into a fixed number of sectors (also called districts, zones, or responsibility areas) according to criteria and constraints regarding the geographic, demographic, political or organizational characteristics of the territory and planning purposes [5, 7]. The basic units can be points

or can be small geographic areas. This process usually aims at better organizing or simplifying a large problem into smaller sub-problems.

The most common criteria in sectorization are Equilibrium (or balance), Compactness, and Contiguity. Other criteria include limited capacity and desirability [16].

Sectorization problems appear in several areas of actuation. Classical applications are political districting [13, 19], the design of school zones [2], municipal solid waste collection areas [10], health care management [1] and Commercial Territory Design [15].

This paper deals with a particular Sectorization Problem: Commercial Territory Design (CTD). This application divides the market or operation area into disjoint sub-territories or responsibility regions according to certain planned criteria. Commonly, these criteria involve demand and the capacity response of the company. Several examples are presented in the literature. In [17], a bi-objective combinatorial model is presented, considering (1) the minimization of the dispersion on each territory (or sector) and (2) the minimum of the maximum relative deviation between the number of customers in each territory and the target value (the mean). The authors propose an improved scatter search-based framework due to the large instances addressed in their work. In [15], the authors considered territory compactness, territory balance and territory connectivity as the planning criteria to satisfy in a real-world application related to the bottled beverage distribution industry. The incompatibility requirements and the similarity to a previous plan make the problem more difficult. An original mixed-integer linear programming model that covers all these characteristics and a strategy to reduce the size of the problem and the search space was developed. A basic model using common characteristics to most applications appears in [6]. This paper introduced a general and interactive framework with GIS to help decision-makers solve different applications. A solution algorithm based on the GRASP approach is presented in [14]. Territorial contiguity, connectivity, compactness, and balanced workload were the criteria considered in a problem motivated by a real-world application. A districting problem related to on-time last-mile delivery with quality service is described in [18]. This quality service is measured as the percentage of on-time deliveries. A two-stage approach was considered: the first, more stable, is the districting planning and the second, at an operational level, is related to last-mile routing delivery decisions. Routing is often associated with CTD and is used to evaluate the quality of the sectors.

When the instances of sectorization problems change over time, whether fluctuating demand or new locations that appear or old locations not needed, the problem is known as dynamic sectorization [8]. When facing this change, the defined sectors must adapt accordingly, and solutions for several periods are necessary. Dynamic problems are more difficult to solve but relevant in a constantly changing world. Interesting examples are the dynamic redesign of the airspace to accommodate changes in traffic demands [20] and the service capacity of public schools considering changes in the supply-demand relationship [3]. Sectorization problems may involve the dynamic concept and the route definition over a discrete planning horizon [9]. The concept of dynamic sectorization is important in the context of this paper, namely because it offers an important line for future work.

### 3 Parcel Delivery Services Case Study

The knowledge generated by the behavior of variables and the relationships between them in real problems allow for evaluating future situations through simulation. Data collection and processing are of great importance in gathering that knowledge. In this case study, real data from a delivery services company was collected, and a simulation was performed to deal better with future delivery services.

The case study addressed in this work relates to a company that delivers parcels in the region of Porto. According to the daily demand, teams are assigned a set of postal codes to work. Assigning the right team to each area is an advantage for the company. Knowing particularly well the working region turns the delivery service more efficient, as it is easier to park and faster to find the destiny door. Hence, the objective of this case study was to define the best area for each parcel delivery team, achieving compact regions for each team that are balanced between each other and sufficiently stable considering the flexibility of the daily demand.

The region of Porto is well identified to the level of the 7-digit postal codes. Daily data of the demand in each of the 7910 postal codes, from a week in May 2022, was given to analyze. The company assured us this was a regular week and was representative of the usual delivery service. The provided sample enclosed many postal codes and parcels, but since the deliveries are only on weekdays, the sample only has five days. Therefore, it was necessary to have more data to understand the service fully. For that reason, a simulation on top of the provided data from the company was necessary.

Table 1 briefly summarises the demand for each 7-digit postal code. Note that one 7-digit postal code does not correspond to just one delivery point; each postal code has many delivery points. The daily demand for a postal code is the number of address points where delivery is due that day. The coordinates of a central point in the 7-digit postal code are representative of all the parcels to be delivered in that postal code. Note that, usually, a 7-digit postal code has many address points, but they are close to each other in urban areas; they usually comprehend just a small part of one street. The average daily demand and the standard deviation are shown, along with the maximum demand observed in a postal code, total demand and the number of postal codes with zero demand daily. For example, on Monday, each postal code

**Table 1** Sample statistics of the daily demand in each 7-digit Postal Code

|                 | Monday | Tuesday | Wednesday | Thursday | Friday |
|-----------------|--------|---------|-----------|----------|--------|
| Average         | 1.01   | 1.27    | 1.19      | 1.33     | 1.23   |
| Std.Deviation   | 6.92   | 7.57    | 7.45      | 8.32     | 9.75   |
| Nb. of Zeros    | 5105   | 4429    | 4484      | 4214     | 4653   |
| 99th Percentile | 14     | 17      | 14        | 14       | 13     |
| Max.            | 554    | 579     | 582       | 587      | 660    |
| Total demand    | 7967   | 10062   | 9398      | 10495    | 9746   |

**Table 2** Postal Codes categories

| Postal Codes Categories<br>(number of zero demand days) | Number of<br>Postal Codes | Number of Postal Codes<br>with constant demand |
|---------------------------------------------------------|---------------------------|------------------------------------------------|
| 0                                                       | 647                       | 14                                             |
| 1                                                       | 702                       | 86                                             |
| 2                                                       | 1096                      | 353                                            |
| 3                                                       | 1869                      | 1170                                           |
| 4                                                       | 3596                      | 3596                                           |
| Total                                                   | 7910                      | 5219 (65.98%)                                  |

has a demand that varies between 0 and 554 delivery points, with a total of 7967. The mean is 1.01 delivery points, and the standard deviation of 6.92. The small value of the mean is due to the large number of postal codes without deliveries (5105 postal codes with zero demand). In any day of the week, the minimum was 0, the median demand was 1, and the third quartile was 4 delivery points. The latter states that 75% of the postal codes have at most 4 delivery points daily. The 99th percentile differs each day but is still a very low value, as shown in Table 1.

This week, a very heterogeneous demand was found in the postal codes analyzed. While some postal codes had constant demand over the week, others had fluctuating demand from one day to the next. Table 2 shows the distribution of postal codes according to the number of days in which they have zero demand and according to the number of days with constant demand. Most postal codes ( $\approx 66\%$ ) have constant demand on all required days (non-zero demand days), and the remaining 34% have different demands on the required days. There are only 647 postal codes that require deliveries on all days of the week, from which only 14 had constant demand all week, and therefore the remaining 633 postal codes have variations in the demand during the week. Also, 702 postal codes have demand on all days of the week except one, from which 86 postal codes have the same demand in each of those 4 days of the week, and the remaining 616 postal codes have differences in the demand of the 4 required days of the week. 1096 postal codes have demand only on three days of the week (2 zero demand days), 1869 have demand on two days (3 zero demand days), and 3596 postal codes require deliveries only on one day of the week (4 zero demand days).

Due to confidentiality reasons, we cannot share much more detail about the daily demand.

A Monte Carlo simulation process for building dynamic data involved two steps in analyzing a larger data set. The frequency and intensity of deliveries in the provided five days of the week are used to build a sample for the five days of the following week. It was considered that deliveries are independent of the day-to-day.

**Step 1: use a binary variable to decide if a postal code will have demand on a certain day**

In the provided sample, not all postal codes have demand every day. To simulate this characteristic, in the first step, it is decided, for each postal code, on each day, if that location is going to be discarded on that day (**No**: *Postal code has no demand on that day*) or if that location will have parcels to be delivered there (**Yes**: *Postal code has some demand on that day*). The demand in the previous five days is analyzed to decide between *Yes* and *No* for this variable. If that postal code had demand in  $k$  days of the previous five days ( $k = 1, 2, \dots, 5$ ), then it is considered that the probability of having demand in any of the coming days is  $k/5$ . For each day of the following week and each postal code, a binary value is assigned to it, using a Bernoulli distribution  $B(1, \theta)$  with probability  $\theta = k/5$ . For each postal code, the expected mean number of days with service in the following week is, therefore  $k$ , with a standard deviation of  $\sqrt{n\theta \times (1 - \theta)} = \sqrt{k(1 - k/5)}$ .

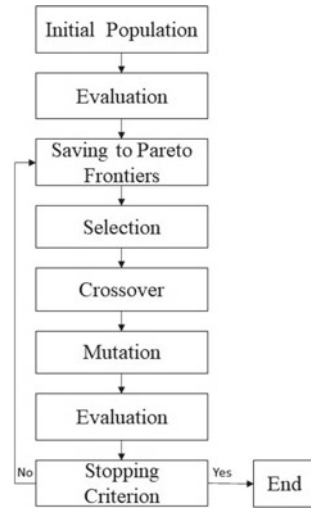
**Step 2: use a quantitative variable to decide the demand of a postal code on a certain day**

If in step 1, it is decided that a postal code will have service on a given day, then it is necessary to assign a certain **quantity** of parcels to be delivered. This demand quantity is simulated with a Normal distribution  $N(\mu, \sigma)$ . To estimate  $\mu$  and  $\sigma$ , the sample mean and standard deviation of the number of parcels in that postal code are computed, but only regarding the days when demand existed in that location. If the deliveries on those days are constant, the standard deviation observed is null, and the value remains constant on the simulated day. The zero demand days are not considered in the computation of these statistics. f

## 4 Results

The sectorization of the postal codes in Porto was solved using a multi-objective optimization (MOO) method called NSGA-II proposed by [4]. As with all MOO methods, NSGA-II evaluates the solutions for several objectives separately and simultaneously and provides a Pareto optimal set. An extensive review of NSGA II, its extensions, and applications is presented in [21]. This method was chosen because it is known to obtain good results with up to three objectives, (see [23]). That evaluation uses two parameters: (i) domination count, and (ii) set of domination members. The dominance count indicates how many other solutions dominate a certain solution. A set of dominating members contains the solutions that a certain solution dominates. After creating these two parameters for each solution, the ones with the lowest domination count are located in the first Pareto frontier. The ones dominated just by the solutions on the first Pareto frontier compose the second Pareto frontier. The ones dominated by just the first and second frontiers constitute the third Pareto frontier. Allocation of all the solutions in the Pareto frontiers continues by this logic. The solutions in the earlier Pareto frontiers are considered better than those in the latter. Moreover, the solutions in the same Pareto frontier are assumed to be equally good. NSGA-II proposes a concept called Crowding Distance (CD) to select between the solutions located in the same Pareto frontier. CD represents the concentration of the solutions

**Fig. 1** The NSGA-II framework



in the Pareto frontier; a solution with a lower value of CD is more crowded with other solutions. Solutions with higher CD are assumed to be better and more representative of the population to increase diversity.

It is possible to see the basic framework of NSGA-II in Fig. 1. As is seen, one full iteration in NSGA-II includes evaluation, selection, crossover, and mutation. The procedure mentioned above (i.e., forming the Pareto frontiers and CD) constitutes the evaluation step. Selection is necessary to conduct crossover, the step where two solutions are mated to generate offspring solutions. In this paper, we used tournament selection. Tournament selection picks two random solutions, evaluates their goodness according to the Pareto frontiers in which they are located or according to CD (in case they are in the same Pareto frontier), and selects the one with a better performance. This procedure is repeated two times to select two solutions as parents. Afterwards, the selected solutions are mated in the crossover step. We used multipoint crossover, a proper method for the genetic encoding system used in this work. The matrix-based genetic encoding system is designed to represent best the sectorization problems, where columns stand for sectors and rows are the basic units (i.e., postal codes in Porto). Each row of the binary matrix takes the value of one where the basic unit is assigned to a sector, and zero otherwise. This encoding system has two restrictions: (i) a basic unit can be assigned only to one sector, and (ii) each sector must have at least one basic unit. Finally, we used random mutation concerning a defined probability level. The sector of a basic unit changes if the mutation probability holds. We used a defined number of generations as the stopping criterion.

The algorithm was written in Python 3.7 and the method is executed on a PC Intel Core i7- 8550U at 1.8 GHz and Win X64 operating system. The parameters used to execute NSGA-II in each period were 500 generations, a population of size 200 and a mutation rate of 0.01. Considering the results found by [11], these parameters were

**Table 3** Number of postal codes with expected demand in the simulation data set

| Day | Basic units |
|-----|-------------|
| 6   | 3292        |
| 7   | 3311        |
| 8   | 3322        |
| 9   | 3320        |
| 10  | 3344        |

selected. The authors conducted multiple test runs with different population sizes and mutation probabilities, demonstrating that larger populations and smaller mutation probabilities outperformed other configurations in three performance metrics relevant to MOO, namely the error ratio, inverted generational distance, and spacing. Therefore, a high population size and low mutation probability were employed. Furthermore, the number of generations was set to 500. As per [22], using a large population compensates for the number of generations in evolutionary algorithms. Consequently, 500 generations proved to be sufficient for obtaining desirable solutions. In Table 3, the number of basic units represents the number of postal codes with expected daily demand. These units and their demand are simulated from the five days sample of real data with 7910 total units, as explained in Sect. 3. The number of sectors considered was 100.

As mentioned in Sect. 2, the previous studies considered three objectives, namely, equilibrium, compactness and contiguity as the most relevant CTD problems [15, 17]. Since, in our data the information regarding the connections between the basic units is unavailable, the two objectives are considered, namely, equilibrium and compactness. The definitions for these objectives are adopted from [12, 16] and represented in Eqs. (1) and (2).

Equilibrium means the balance between sectors in terms of total demand [16].

$$std'_{B_q} = \sqrt{\frac{1}{K-1} \sum_{j=1}^K (q_j - \overline{B_q})^2} \text{ with } \overline{B_q} = \frac{\sum_{j=1}^K q_j}{K}, \quad q_j = \sum_i x_{ij} \times d_i, \forall j \quad (1)$$

Here,  $q_j$  is the total demand in sector  $j$ , where  $x_{ij}$  is 1 if the basic unit  $i$  with demand  $d_i$  is located in sector  $j$ .  $\overline{B_q}$  represents the mean demand and  $K$  is the number of sectors. This value is used to compute the standard deviation ( $std'_{B_q}$ ) from the mean demand value in each sector. A smaller value of  $std'_{B_q}$  represents a better equilibrium.

Compactness refers to the density of each sector [12]. Equation (2) represents the measure used when the furthest distance is considered to define the compactness.  $C_f$  is the total sum of distances, in all sectors, from the centroid of each sector to its furthest basic unit.

$$C_f = \sum_{j=1}^K dist(o_j, x_{.j}^F) \quad (2)$$

**Table 4** Performance of solutions from Day 6 to Day 10

| Day | Equilibrium $std'_{B_q}$ | Compactness $C_f$ | Computation time |
|-----|--------------------------|-------------------|------------------|
| 6   | 126.14                   | 2.15              | 2:13:24          |
| 7   | 113.94                   | 2.30              | 1:54:51          |
| 8   | 97.61                    | 2.79              | 1:55:30          |
| 9   | 107.86                   | 2.51              | 1:54:42          |
| 10  | 95.13                    | 3.37              | 1:56:06          |

In Equation (2),  $o_j$  shows the mass centroid of the sector  $j$ , and  $x_j^F$  is the basic unit belonging to sector  $j$  which is furthest from its centroid  $o_j$ . The smaller the  $C_f$  values, the better the compactness of the sectorization.

The algorithm was run for the 5-day sample with the simulated data considering the minimisation of these two objectives. The numerical performance of each solution is represented in Table 4. The first two columns show the performance of the solutions for equilibrium and compactness, respectively. Moreover, the last columns represent the computation time (presented in hours) the algorithm took to run the given instance.

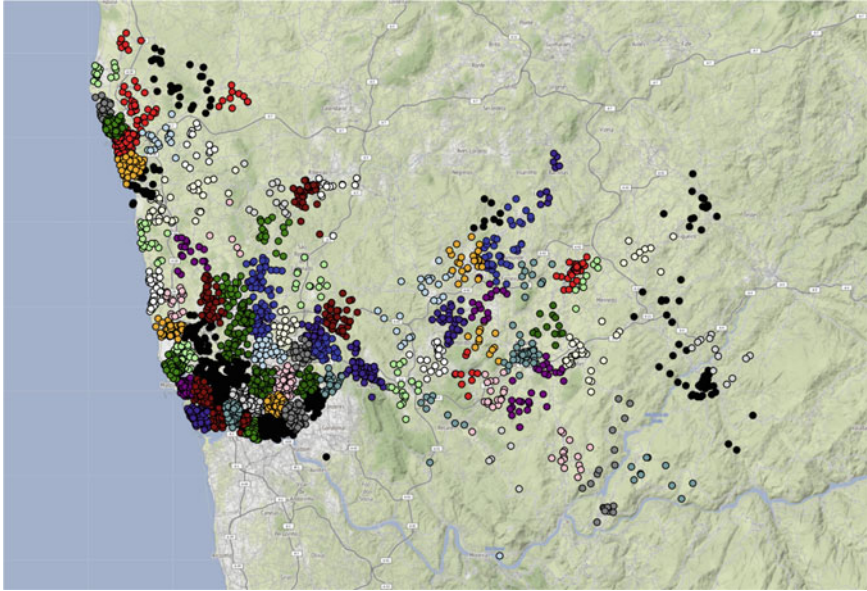
The solutions for the sectorization of the parcel delivery service for each day are presented in Figs. 2, 3 and 4.<sup>1</sup> Different colors correspond to postal codes that have been assigned to different sectors. Everyday, the basic units with expected demand (Table 3) have been assigned to 100 sectors.

Each point in the map of the Porto region corresponds to the location of a representative point of the 7-digit postal code, which contains not one, but several delivery addresses/points. The demand of each postal code is the number of delivery points expected for that postal code.

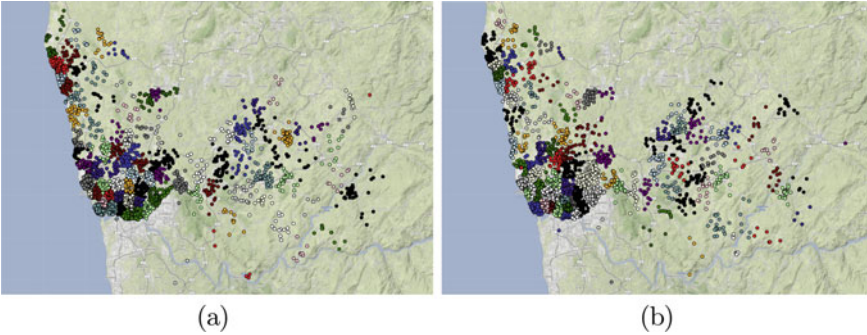
## 5 Conclusions and Future Work

This paper considered a real Parcel Delivery Services problem. The approach followed was based on sectorization, which is quite appropriate in such a context. We resorted to simulation to generate new and more extensive data. The multi-objective algorithm NSGAII was used to solve the problems. The simulation considered the existing knowledge about the periodicity and the quantities of the deliveries in each of the 7910 7-digit postal codes. A full week (five days) was simulated. For each day, 100 sectors were produced. Each sector represents an area of responsibility for the delivery truck. The two objectives considered were Equilibrium, because the demand must be similar between sectors, and Compactness, due to the importance of minimizing the total space dispersion of the sectors. Visualizing the results on the map with different colors makes it easy for the company to understand the geographic

<sup>1</sup> These figures can be viewed on a larger scale [here](#).



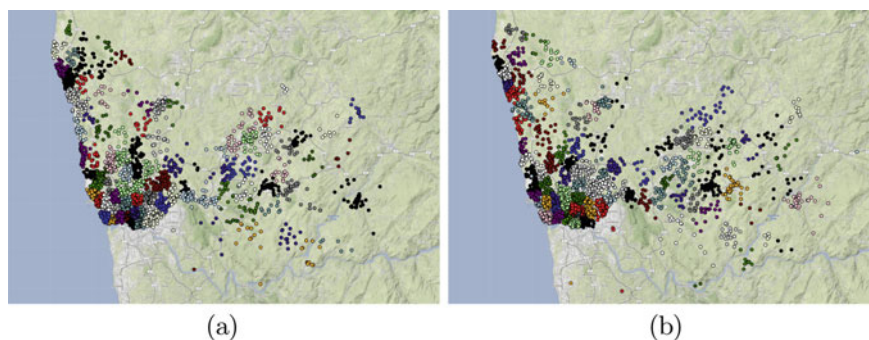
**Fig. 2** Sectorization of parcel delivery services for day 6, Monday



**Fig. 3** Sectorization of parcel delivery services for (a) day 7, Tuesday, and (b) day 8, Wednesday

planning of the deliveries. Due to simulation, new scenarios can be drawn and readily discussed by decision-makers. Considering only five values for each postal code may seem a study limitation. However, it should be noted that each postal code is not just a delivery point, since it may represent the accumulation of many hundreds of points/clients.

In future work, a robustness measure can be considered to incorporate similarity in the response from one day to another. Progressing to another level of approach, treating this case of sectorization as Dynamic Optimization (instead of just dynamic instances) would contribute to the robustness of the solutions. The optimization will



**Fig. 4** Sectorization of parcel delivery services for (a) day 9, Thursday, and (b) day 10, Friday

consider all the connected periods (days) in an integrated way. Sectors should be similar to day-to-day. To accommodate the desired continuity in the response, the solution of day  $i$  could be used as a seed to start the algorithm for the day  $i + 1$ , maintaining stability among working teams, at least in the postal codes with daily demand.

**Acknowledgements** This work is financed by the ERDF–European Regional Development Fund through the Operational Programme for Competitiveness and Internationalization–COMPETE 2020 Programme and by National Funds through the Portuguese funding agency, FCT–Fundação para a Ciência e a Tecnologia within projects POCI-01-0145-FEDER-031671 and UIDB/05422/2020.

## References

1. Benzarti, E., Sahin, E., Dallery, Y.: Operations management applied to home care services: analysis of the districting problem. *Decis. Support Syst.* **55**, 587–598 (2013)
2. Bouzarth, E.L., Forrester, R., Hutson, K.R., Reddoch, L.: Assigning students to schools to minimize both transportation costs and socioeconomic variation between schools. *Socioecon. Plann. Sci.* **64**, 1–8 (2018)
3. Cui, H., Wu, L., Hu, S., Lu, R.: Measuring the service capacity of public facilities based on a dynamic voronoi diagram. *Remote Sens.* **13**(5), 1–15 (2021)
4. Deb, K., Pratap, A., Agarwal, S., Meyarivan, T.: A fast and elitist multiobjective genetic algorithm: NSGA-II. *IEEE Trans. Evol. Comput.* **6**(2), 182–197 (2002)
5. Kalcsics, J., Nickel, S., Schröder, M.: Towards a unified territorial design approach — applications, algorithms and GIS integration. *TOP* **13**(1), 1–56 (2005)
6. Kalcsics, J., Nickel, S., Schröder, M.: A generic geometric approach to territory design and districting. *Berichte des Fraunhofer ITWM* **153**, 1–32 (2009)
7. Kalcsics, J., Ríos-Mercado, R.Z.: Districting Problems. In: Laporte, G., Nickel, S., Saldanha da Gama, F. (eds.) *Location Science*. Springer International Publishing, pp. 705–743 (2019)
8. Lei, H., Laporte, G., Liu, Y., Zhang, T.: Dynamic design of sales territories. *Comput. Oper. Res.* **56**, 84–92 (2015)
9. Lei, H., Wang, R., Laporte, G.: Solving a multi-objective dynamic stochastic districting and routing problem with a co-evolutionary algorithm. *Comput. Oper. Res.* **67**, 12–24 (2016)

10. Lin, H.Y., Kao, J.J.: Subregion districting analysis for municipal solid waste collection privatization. *J. Air Waste Manag. Assoc.* **58**(1), 104–111 (2008)
11. Öztürk, E., Rocha, P., Sousa, F., Lima, M., Rodrigues, A.M., Ferreira, J.S., Nunes, A.C., Lopes, C., Oliveira, C.: An application of preference-inspired co-evolutionary algorithm to sectorization. In: *International Conference Innovation in Engineering*, Springer, pp. 257–268 (2022)
12. Öztürk, E.G., Rodrigues, A.M., Ferreira, J.S.: Using AHP to deal with sectorization problems. In: *Proceedings of the International Conference on Industrial Engineering and Operations Management*, pp. 460–471 (2021)
13. Ricca, F., Scozzari, A., Simeone, B.: Political districting: from classical models to recent approaches. *Ann. Oper. Res.* **204**(1), 271–299 (2013)
14. Ríos-Mercado, R.Z., Fernández, E.: A reactive GRASP for a commercial territory design problem with multiple balancing requirements. *Comput. Oper. Res.* **36**(3), 755–776 (2009)
15. Ríos-Mercado, R.Z., López-Pérez, J.F.: Commercial territory design planning with realignment and disjoint assignment requirements. *Omega* **41**(3), 525–535 (2013)
16. Rodrigues, A.M., Ferreira, J.S.: Measures in sectorization problems. In: *Operations Research and Big Data*. Springer, pp. 203–211 (2015)
17. Salazar-Aguilar, M.A., Ríos-Mercado, R.Z., González-Velarde, J.L., Molina, J.: Multiobjective scatter search for a commercial territory design problem. *Ann. Oper. Res.* **199**(1), 343–360 (2012)
18. Sandoval, M.G., Álvarez-Miranda, E., Pereira, J., Ríos-Mercado, R.Z., Díaz, J.A.: A novel districting design approach for on-time last-mile delivery: an application on an express postal company. *Omega (United Kingdom)* **113**, 102687 (2022)
19. Swamy, R., King, D.M., Jacobson, S.H.: Multiobjective Optimization for Politically Fair Districting: A Scalable Multilevel Approach. *Oper. Res.* **71**(2), 536–562 (2022)
20. Tang, J., Alam, S., Lokan, C., Abbass, H.A.: A multi-objective approach for dynamic airspace sectorization using agent based and geometric models. *Transp. Res. Part C* **21**, 89–121 (2012)
21. Verma, S., Pant, M., Snasel, V.: A comprehensive review on NSGA-II for multi-objective combinatorial optimization problems. *IEEE Access* **9**, 57757–57791 (2021)
22. Vrajitoru, D.: Large population or many generations for genetic algorithms - implications in information retrieval. In: *Soft Computing in Information Retrieval*, pp. 199–222. Springer (2000)
23. Zheng, W., Tan, Y., Fang, X., Li, S.: An improved MOEA/D with optimal DE schemes for many-objective optimization problems. *Algorithms* **10**(3) (2017)

# Decision Support System for Scheduling the Production of Screw Caps in a Flexible Job-Shop Environment



José Filipe Ferreira and Rui Borges Lopes

**Abstract** This work presents a decision support system developed in a company that produces customizable capsules for bottles. More specifically, in the production department of wine screw caps, whose production is of the flexible job-shop type. Due to the high growth of the company's department under analysis, both in terms of the number of orders and resources, production scheduling has become increasingly complex. This complexity led to the need to develop a decision support system, based on a constraint programming model, which would allow considering the main characteristics of the production environment. The proposed model was implemented in the CP Optimizer, a component of CPLEX, and two different objective functions were tested: minimization of makespan and minimization of maximum tardiness. This model was integrated with external tools, to receive inputs and provide data in a format that is more easily interpreted by the decision maker, using Gantt charts. The model was validated with real instances, obtained from the company's historical data, achieving a predicted reduction of around 19% in the number of orders delayed, when comparing with the actual scheduling carried out.

**Keywords** Decision support system · Scheduling · Flexible job-shop · Bottle cap industry

## 1 Introduction

The work detailed in this paper took place in a Portuguese company which will be named Alpha for confidentiality reasons. Alpha was founded about 100 years ago and is a medium-sized company with around 100 employees. The company is in the north of Portugal in a facility with a total area of 50,000 m<sup>2</sup>, from which 15,000 m<sup>2</sup>

---

J. F. Ferreira (✉)

DEGEIT, University of Aveiro, Campus Universitário de Santiago, 3810-193 Aveiro, Portugal  
e-mail: [jfilipeferreira@ua.pt](mailto:jfilipeferreira@ua.pt)

R. B. Lopes

CIDMA/DEGEIT, University of Aveiro, Campus Universitário de Santiago,  
3810-193 Aveiro, Portugal  
e-mail: [rui.borges@ua.pt](mailto:rui.borges@ua.pt)

are covered areas. It produces bottle capsules, mostly for the wine industry, and its production is composed of several types of capsules: PVC, poly laminate, aluminium, tin, hoods, and screw caps.

The aluminium screw caps production sector will be the focus of this work. Produced in aluminium, these capsules are intended for bottles containing wine, drinks with a high degree of alcohol and olive oil. These capsules can be customized according to customers requests, leading to several customizations, which often implies variable production processes from order to order. This sector was the one that grew the most in Alpha and, therefore, drawing the biggest investment. Investment was mostly aimed at enabling a wide range of product customization options, in order to gain competitive advantage. Consequently, the increase in number of orders and machines led to an increasingly bigger scheduling problem—each capsule can have up to ten operations and several machines are available for each of these operations. While the scheduling has been done based on the experience and intuition of the responsible of the sector, the increase in size and complexity makes this task increasingly harder, arising the need for a decision support system (DSS) to support this task.

Correct implementation of a DSS often requires a deep understanding of the problem, identification of the main target users and use cases, identification of appropriate technologies, and an implementation and follow up that fits the needs of target users. In this case, the target user is the decision maker, the responsible of the production sector. This user has higher education, a good knowledge of the underlying problem (although not necessarily of the used OR methods), reasonable computer literacy, and is to use the DSS infrequently (usually, once or twice a week). With this profile, from an interaction and information visualization point of view, the main concern is that the system be results oriented, providing easy and efficient access to solutions and relevant data.

The main purpose of this work is to present a DSS, developed in a Portuguese screw cap manufacturing company, with the goal of helping the scheduling of the aluminium screw caps production sector. The DSS is to import data from Alpha's production planning and control software, find the best scheduling solution, and present in a way that is easily understandable by the decision maker.

This paper is structured as follow. This section introduced the main motivations and objectives of the work. Section 2 introduces some concepts concerning scheduling problems and briefly reviews models and applications relevant to the case under analysis. In Sect. 3 the production process is detailed, allowing to characterize the underlying problem. Section 4 presents the DSS framework, its main components and software used. Results are presented in Sect. 5, where different objective functions are tested and a comparison is made to the company's actual scheduling. Finally, conclusions and future research directions are drawn in Sect. 6.

## 2 Scheduling Problem and Its Applications

The goal of the scheduling problem is to find the best order to process a given number of jobs on a specified number of resources able to execute them over given time periods [1]. Characterized by a virtually unlimited number of problem types [2], it is in manufacturing industries that scheduling activities gain particular importance. Scheduling is a key factor for productivity, as it can improve on-time delivery, reduce inventory, cut lead times, and improve the utilization of bottleneck resources [3]. However, these problems are frequently difficult to solve, especially in the limited decision time often available.

The term job-shop is used in production environments making highly variable products, usually in small quantities, where each job often has its own sequence of operations (called routing). According to Sonmez and Baykasoglu [4], the job-shop scheduling problem (JSP) belongs to a branch of production programming, which is among the most difficult combinatorial optimization problems. The flexible job-shop scheduling problem (FJSP) is an extension of the JSP which incorporates flexibility in the routing of products, by allowing operations in jobs to be processed in one machine out of a set of alternative machines [5].

As the FJSP is strongly  $\mathcal{NP}$ -hard [6], only in small instances can the optimal be obtained in reasonable time. Up until recently, problems with 20 jobs and 20 machines were still considered very hard to solve to optimality [7].

Recent formulations for the FJSP can be found in [8, 9]. The former is among the most impactful works in the literature, where the authors propose a formulation and use it to solve a set of 20 small and medium sized instances, which have become among the most used for benchmarking. The model by Özgüven et al. [9] was able to obtain better results than Fattahi et al. [8] in terms of computation time and objective function values. Demir and İşleyen [10] further compare Özgüven et al.'s formulation with other well-known models in the literature, considering it to be the most competitive.

Several real-world applications of the FJSP in manufacturing industries can be found in the literature. In the automotive industry, Calleja et al. [11] address a FJSP with transfer batches with the goal of minimizing the average tardiness of production orders. For semiconductor wafers manufacturing, where products follow a re-entrant or re-circulating flow through different machines, in [12] a heuristic approach was developed to minimize the total weighted tardiness. A similar objective was used in [13] for plastic injection mold manufacturing. For the printing and boarding industry, Vilcot et al. [14] propose a multi-objective approach minimizing makespan and maximum tardiness.

In all of these applications it was recognized by practitioners the importance of prioritizing tardiness as an objective, either independently or alongside the commonly addressed objective of minimizing makespan. This observation holds true for the case study presented herein.

### 3 Problem Definition

To develop an appropriate DSS it is important that the underlying problem be fully understood, and its main defining aspects incorporated into the systems' model. To that end, it was necessary to have a clear understanding of the production process, from which the main characteristics were identified. The production process and the problem characterization are detailed in the following subsections.

#### 3.1 *Production Process*

The production process of aluminium screw caps in Alpha requires that products go through up to 10 different processes.

Each production run starts by cutting aluminium plates into smaller sized plates, which then go to a pressing machine. From each small plate, 55 screw caps are obtained in the pressing machine in the second process. Process 3 consists of printing a picture in the top of the product and is not always required. In the following process, called stretching, the product acquires its final shape and is moved in containers to the next stage. Processes 5 and 6 are only to be performed when a drawing or picture must be printed on the side of the screw cap.

Processes 7, 8 and 9 are usually done in the same machine. The screw cap is opened, for it to be applied in bottles later on, and a disc or seal is applied, to ensure the bottle will remain closely sealed. Process 8, embossing, is only performed if the clients require it in the product specifications. Finally, in process 10, the products are stored in cardboard boxes, to later be stacked in pallets and sent to the warehouse. For a more detailed flowchart and description of the typical production process of screw caps, the reader is referred to [15].

Given the characteristics of the production process and shop floor layout, it was possible to group the abovementioned processes into three main groups: cutting and stretching, processes 1 to 4; side printing, processes 5 and 6; and knurling and cutting, processes 7–10. By considering each of the available production lines or machines in these main group of processes as single resources, it was possible to significantly reduce the complexity of the scheduling problem. Figure 1 shows the simplified production process with 3 resources for cutting and stretching, 2 resources for side printing, and 3 resources for knurling and cutting. This will be the problem addressed henceforth.

With this grouping of processes, the size of the problem has been reduced to at most 3 operations for each job, to be performed in 8 machines. Given that the scheduling is done on a weekly basis, with around 25 jobs, it was considered appropriate to use exact methods to solve this problem.

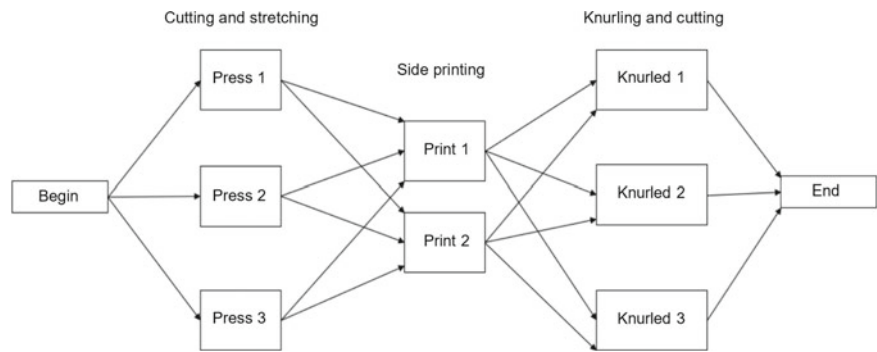


Fig. 1 Flowchart of the simplified production process of screw caps in Alpha

3.2 Problem Characteristics and Formulation

The problem detailed previously has the characteristics of a FJSP, which is described as follows.

There is a set  $J$  of independent jobs,  $J = \{j_1, j_2, \dots, j_n\}$ , where each of the jobs has its own processing order going through a set of machines  $M = \{m_1, m_2, \dots, m_m\}$ . To complete each job  $i$  a number of  $l_i$  ordered operations have to be done.  $O_{i,j}$  is the operation  $j$  of job  $i$ , which can only be performed in one of the machines in a given set  $M_{i,j} \subseteq M$ , with the corresponding processing time denoted  $t_{i,j,k}$  (with  $k$  being the machine where  $O_{i,j}$  is performed). It is assumed that all  $t_{i,j,k}$  are known and fixed, and setup and transportation times are not considered, negligible, or included in  $t_{i,j,k}$ . Moreover, each machine in  $M$  can only perform one operation at a time and pre-emption is not allowed.

The goal is usually to minimize the maximum completion time over all jobs, called makespan, and denoted as  $C_{\max}$ . A constraint programming (CP) formulation for the defined FJSP is the following.

|                            |                                                                                 |
|----------------------------|---------------------------------------------------------------------------------|
| <b>Decision variables:</b> |                                                                                 |
| $op_{i,j,k}$               | Interval object representing the start and end time of $O_{i,j}$ in machine $k$ |
| $C_{\max}$                 | Makespan                                                                        |

**Functions:**

|                           |                                                                                                                                                              |
|---------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <i>endBeforeStart</i> (.) | This function constrains that if predecessor and successor interval variables are present, then the successor cannot start before the end of the predecessor |
| <i>noOverlap</i> (.)      | This function ensures that the interval variables present do not overlap with each other                                                                     |
| <i>alternative</i> (.)    | This function enforces that only one interval from the given set of intervals can be selected                                                                |
| <i>endOf</i> (.)          | If an interval variable is present, this function returns its end time; otherwise, 0 is returned                                                             |

**Constraints:**

$$\text{endBeforeStart}(op_{i,j-1,k}, op_{i,j,k'}) : i \in J, j \in \{2, \dots, l_i\}, k \in M_{i,j}, k' \in M_{i,j} \quad (1)$$

$$\text{noOverlap}(op_{i,j,k}, op_{i,j',k}) : \quad (2)$$

$$i \in J, j \in \{1, \dots, l_i\}, k \in M_{i,j}, j' \in \{1, \dots, l_i\} | j' \neq j$$

$$\text{alternative}(op_{i,j,k} | k \in M_{i,j}) : i \in J, j \in \{1, \dots, l_i\} \quad (3)$$

$$C_{\max} \geq \text{endOf}(op_{i,j,k}) : i \in J, j \in \{1, \dots, l_i\}, k \in M_{i,j} \quad (4)$$

**Objective function:**

$$\text{Minimize } C_{\max} \quad (5)$$

Constraints (1) ensure the correct ordering of operations within each job. Constraint set (2) schedules operations to machines, preventing overlap, while constraints (3) makes that operations can only be processed on one machine from the alternative machine set. The makespan value is captured due to (4).

In the production sector under analysis, number of tardy jobs was considered the main evaluation metric. However, minimization of the maximum delay (i.e. when the difference between due date and completion of a job is negative) was also considered important—this is called maximum tardiness,  $T_{\max}$ . As the company was still exploring different objectives, the latter objective was chosen, as it allowed a more comprehensive evaluation of the job scheduling performance and a more fitting comparison with makespan.

To address this objective, the due date  $d_i$  of each job  $i \in J$  is to be given, and the following variables are introduced:  $T_i$ , tardiness of job  $i \in J$  and  $T_{\max}$ . Constraints (4) are replaced by:

$$T_i = \text{endOf}(op_{i,l_i,k}) - d_i : i \in J \quad (6)$$

$$T_{\max} \geq T_i : i \in J \quad (7)$$

which calculate the tardiness of each job and obtain  $T_{\max}$ , to be minimized.

## 4 Decision Support System

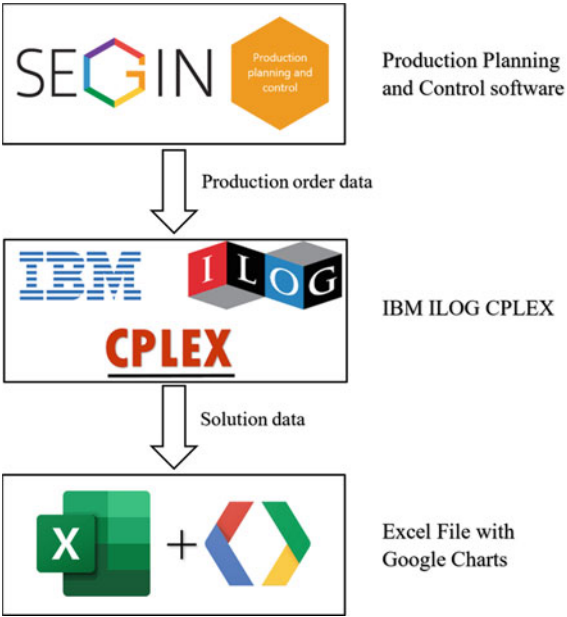
This section shows the adopted framework, and details its main components.

### 4.1 Framework

The proposed DSS uses several software and aims to provide an easy to interpret scheduling solution to the decision maker.

The general framework of the DSS can be seen in Fig. 2. The company’s production and control software (SEGIN) stores all the required data to perform the scheduling. A program was developed to quickly export this data into the IBM ILOG CPLEX tool, where the model was implemented. The model with the imported data is then executed and the results obtained undergo postprocessing, where the solution data is formatted and exported to an Excel file. The Excel file with Google Charts displays the data into a Gantt chart, making the information easily readable by the decision maker. Moreover, by using an Excel file, the data can be effortlessly edited and exported to other tools.

**Fig. 2** General framework of the DSS



## 4.2 Model Implementation and Solution Representation

The mathematical model was implemented in the CP Optimizer, a component of IBM ILOG CPLEX. This tool has been extensively used in scheduling problems for over 20 years, being one of the most successful application areas of CP [16]. Main advantages of this tool are [16]: accessible not only to CP experts but also to software engineers; simple, non-redundant and uses a minimal number of concepts to reduce the learning curve for new users; and provides a robust and efficient search algorithm so that users only need to focus on the model.

The implemented model can be seen in Fig. 3. Besides the standard FJSP formulation, two additional case-specific sets of constraints were added: to handle products with less operations, where the job must occur in the same production line (e.g. if Press 1 is chosen for the first operation, the following operation must be done in Print 1); and to address a production constraint concerning printing the top of screw caps, which is only needed for some products and can only be done in Press 1.

Once the CP optimizer obtains a solution, the data is formatted using postprocessing and sent to an Excel file. In the Excel file the user can easily navigate the data, and possibly display it in a Gantt chart. This graphical representation of the solution was considered the most intuitive.

```

using CP; //https://www.ibm.com/analytics/cplex-cp-optimizer

//Definição de Variáveis
int nEncomendas...;
int nMaq...;
range Encomendas = 1..nEncomendas;
range Maquinas = 1..nMaq;

tuple Operacoes {
  int id; // ID da Operação
  int jobId; // ID do trabalho
  int pos; // Posição no trabalho
};

{Operacoes} Ops = ...;

tuple Alternativas {
  int opId; // ID da Operação
  int maq; // Máquina
  int tp; // Tempo de Processamento
};

{Alternativas} Modes = ...;

// Posição da última operação da encomenda j
int jlast[j in Encomendas] = max(o in Ops: o.jobId==j) o.pos;

//int DataEntrega[Encomendas]=...; //Dados de entrega (Modelo B)

//Variáveis de decisão:
dvar interval ops [Ops];
dvar interval modes[md in Modes] optional size md.tp;
dvar sequence maqs[m in Maquinas] in all(md in Modes: md.maq == m) modes[md];

//Função Objetivo (função objetivo do modelo A):
minimize max(j in Encomendas, o in Ops: o.pos==jlast[j]) endOf(ops[o]);

//minimize max(j in Encomendas, o in Ops: o.pos==jlast[j]) (endOf(ops[o]) -
DataEntrega[j]); //função objetivo do modelo B)

//Restrições:
subject to {
  // Restrição de precedência entre operações consecutivas o1,o2 de uma encomenda
  forall (j in Encomendas, o1 in Ops, o2 in Ops: o1.jobId==j && o2.jobId==j &&
o2.pos==o1.pos)
    c1:
      endBeforeStart(ops[o1],ops[o2]);

  // Máquinas alternativas para uma determinada operação o
  forall (o in Ops)
    c2:
      alternative(ops[o], all(md in Modes: md.opId==o.id) modes[md]);

  // Operações na máquina m não podem sobrepor-se/overlap
  forall (m in Maquinas)
    c3:
      noOverlap(maqs[m]);

  //Restrição para trabalhos pequenos (realizados de modo contínuo na mesma
  linha)
  forall(md1 in Modes, md2 in Modes: md2.opId==md1.opId &&
md2.mch==md1.maq)
    c4:
      presenceOf(modes[md1])==presenceOf(modes[md2]);

  // Restrição para máquina do Topo (Linha 1) e pequenos trabalhos
  forall (j in Encomendas, o1 in Ops, o2 in Ops: o1.jobId==j && o2.jobId==j &&
o2.pos==o1.pos)
    c5:
      startAtStart(ops[o1],ops[o2]);
}

```

**Fig. 3** Model implemented in CP optimizer. Decision variables and parameters on the left, objective functions and constraints on the right

**Table 1** Overall results of the implemented model on benchmark instances

| Instance set   | Instances | Average time (s) | Average gap (%) | Median gap (%) | Maximum gap (%) |
|----------------|-----------|------------------|-----------------|----------------|-----------------|
| Brandimarte    | 6         | 322.26           | 18.6            | 7.0            | 72.7            |
| Fattahi et al. | 20        | 1.51             | 11.3            | 11.5           | 32.6            |
| Kacem et al.   | 4         | 25.68            | 0               | 0              | 0               |

## 5 Results

In this section the results obtained by the model will be detailed, all obtained using IBM ILOG CPLEX Optimization Studio 12.7, running on an Intel Core 2 Duo E8400 CPU @ 3.00 GHz.

The first step was to test the developed model in the instances from the literature. To that effect, before implementing problem specific constraints, the model was tested in the well-known benchmark instances by Brandimarte [17], Fattahi et al. [8], and Kacem et al. [18]. The overall results, regarding gaps to lower bounds, can be seen in Table 1; the interested reader is referred to [19] for a full breakdown of the results. All small instances were solved to optimality, within reasonable computing times, and results concerning medium-sized instances can be considered competitive.

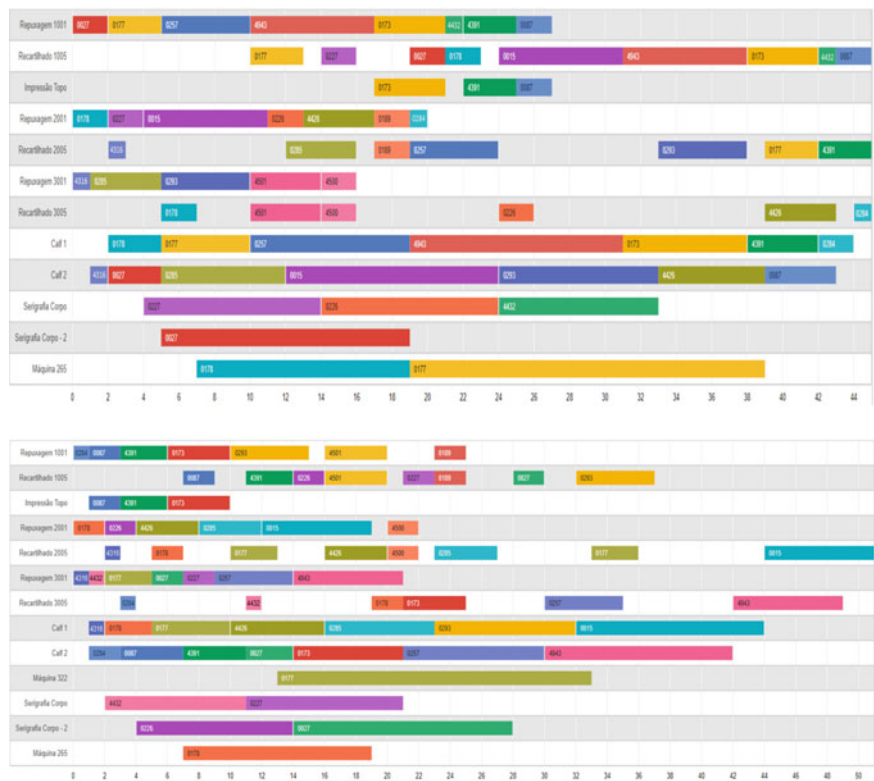
As the model proved competitive in solving the underlying scheduling problem, problem-specific constraints were added, and the model was tested using real data from Alpha. Two newly created instances were based on historical data from two weeks in February 2020, which were considered representative of typical working weeks. The first instance has a total of 20 jobs to be performed, while the second instance has 26 jobs.

For the decision maker, the main goal was often the delivery of customers' orders on time, and the first analysis concerned which objective function would be the most fitting—the first instance was used for this. Afterwards, using the data from the second instance, the obtained solution is to be compared with the scheduling that actually took place, to understand the advantages and disadvantages of using the proposed DSS.

### 5.1 Testing Objective Functions

The typical FJSP focuses on makespan minimization, however, for the production sector under analysis, complying with due dates was highly valued, prompting the test of minimization of maximum tardiness.

Using the model with the makespan minimization objective, the optimal solution was reached in 2.41 s, which will be called Solution A. For maximum tardiness minimization, the model considers a limit on the search for solutions—it stops after



**Fig. 4** Graphical representation of Solution A (top) and Solution B (bottom)

5 million unsuccessful attempts to find a feasible or improved solution. The limit was set to ensure comparable computation times for both models, and to facilitate the prompt delivery of solutions by the DSS. For this instance, it took 39.91 s to obtain a solution, Solution B, with 81% gap. A graphical representation of the solutions can be found in Fig. 4.

Concerning makespan, Solution B takes 6 more time units to complete, representing a 13% gap. However, Solution A has 45% of the jobs delayed, compared with 20% delayed jobs in Solution B. Despite not being the optimal solution for maximum tardiness, Solution B was considered the most interesting from the decision maker's point of view.

## 5.2 *Comparing with Actual Scheduling*

The second newly created instance was used to compare the solution provided by the DSS with the scheduling that took place in the shop floor; both will be called respectively Solution C and Solution R.

The makespan of Solution R is 73 while the makespan of Solution C is 62, representing a reduction of 15%. Concerning the number of delayed jobs, 7 out of the 26 jobs are delayed in Solution R (26.9%), while in Solution C only 2 jobs are delayed (7.7%). Therefore, if the newly proposed scheduling was to be implemented, 5 (19.2%) less orders would be delayed. Of note is that both solutions start essentially with the same jobs, however, they rapidly diverge, due to the number of alternative machines available. For the graphical representation of Solutions C and R, the reader is referred to [19].

## 6 Conclusions

Alpha was faced with a challenge regarding the scheduling of one of its production sectors, as it was increasing in size (number of machines and production lines) and number of different products. The company had no system to help the responsible of the sector in this task, becoming increasingly arduous and less efficient to make the scheduling based purely on his experience and intuition.

To address this a DSS was developed. The system was to communicate with the current software of the company and obtain, with relative ease, a good solution for the scheduling problem. Moreover, it was important that the system would present the obtained solutions in a way that was easy to interpret.

Despite several machines and production lines available for the different operations, it was possible to reduce the size of the underlying problem without compromising its realism. This allowed to employ tools that would enable to reach optimum solutions and leave room for further developments regarding case-specific constraints. The chosen tool was the CP Optimizer, a component of CPLEX, where the mathematical model was implemented. To easily edit and visualize the solution data in a Gantt chart, Excel is used.

Results suggest that the most appropriate objective for this sector would be minimization of maximum tardiness, and that using the proposed DSS would enable a reduction of about 19% in the number of orders delayed. It is worth noting that some aspects valued by the decision maker, such as adaptability to disruptions (e.g. delays, absenteeism, and unscheduled stoppages) and prioritizing skilled workers for demanding operations, have not been addressed yet.

A limitation of this work is that conclusions are based on results of only two weeks, even though considered representative of typical working weeks. Additionally, the scheduling approach does not comprehensively account for the rolling nature, such

as effectively managing jobs already in progress at the start of the week. One possible solution to address this issue is by utilizing idle machines.

Future work should therefore focus on testing on more instances, covering a longer period of time. Moreover, further integration of the DSS with the companies' software may be beneficial, as well as presenting to the user other performance indicators. The implemented model could be improved by trying to incorporate some resilience to disruptions.

**Acknowledgements** This work was supported by The Center for Research and Development in Mathematics and Applications (CIDMA) through the Portuguese Foundation for Science and Technology (FCT—Fundação para a Ciência e a Tecnologia), references UIDB/04106/2020 and UIDP/04106/2020.

## References

1. Pinedo, M.L.: *Scheduling: Theory, Algorithms, and Systems*. Springer, New York, USA (2012)
2. Brucker, P.: *Scheduling Algorithms*, 5th edn. Springer, New York, USA (2007)
3. Tang, L., Liu, G.: A mathematical programming model and solution for scheduling production orders in Shanghai Baoshan Iron and Steel Complex. *Eur. J. Oper. Res.* **182**(3), 1453–1468 (2007). <https://doi.org/10.1016/j.ejor.2006.09.090>
4. Sonmez, A.I., Baykasoglu, A.: A new dynamic programming formulation of  $(n \times m)$  flowshop sequencing problems with due dates. *Int. J. Prod. Res.* **36**(8), 2269–2283 (1998). <https://doi.org/10.1080/002075498192896>
5. Chan, F.T.S., Wong, T.C., Chan, L.Y.: Flexible job-shop scheduling problem under resource constraints. *Int. J. Prod. Res.* **44**(11), 2071–2089 (2006). <https://doi.org/10.1080/00207540500386012>
6. Garey, M.R., Johnson, D.S., Sethi, R.: The complexity of flow shop and job shop scheduling. *Math. Oper. Res.* **1**(2), 117–129 (1976)
7. Zhang, C.Y., Li, P.G., Rao, Y.Q., Guan, Z.L.: A very fast TS/SA algorithm for the job shop scheduling problem. *Comput. Oper. Res.* **35**(1), 282–294 (2008). <https://doi.org/10.1016/j.cor.2006.02.024>
8. Fattahi, P., Saidi-Mehrabad, M., Jolai, F.: Mathematical modeling and heuristic approaches to flexible job shop scheduling problems. *J. Intell. Manuf.* **18**, 331–342 (2007). <https://doi.org/10.1007/s10845-007-0026-8>
9. Özgüven, C., Özbakır, L., Yavuz, Y.: Mathematical models for job-shop scheduling problems with routing and process plan flexibility. *Appl. Math. Model.* **34**(6), 1539–1548 (2010). <https://doi.org/10.1016/j.apm.2009.09.002>
10. Demir, Y., İşleyen, K.S.: Evaluation of mathematical models for flexible job-shop scheduling problems. *Appl. Math. Model.* **37**(3), 977–988 (2013). <https://doi.org/10.1016/j.apm.2012.03.020>
11. Calleja, G., Pastor, R.: A dispatching algorithm for flexible job-shop scheduling with transfer batches: an industrial application. *Prod. Plan. Control* **25**(2), 93–109 (2014). <https://doi.org/10.1080/09537287.2013.782846>
12. Mason, S.J., Fowler, J.W., Carlyle, W.M.: A modified shifting bottleneck heuristic for minimizing total weighted tardiness in complex job shops. *J. Sched.* **5**(3), 247–262 (2002). <https://doi.org/10.1002/jos.102>
13. Caballero-Villalobos, J.P., Mejía-Delgadillo, G.E., García-Cáceres, R.G.: Scheduling of complex manufacturing systems with Petri nets and genetic algorithms: a case on plastic injection

- moulds. *Int. J. Adv. Manuf. Technol.* **69**(9–12), 2773–2786 (2013). <https://doi.org/10.1007/s00170-013-5175-7>
14. Vilcot, G., Billaut, J.: A tabu search algorithm for solving a multicriteria flexible job shop scheduling problem. *Int. J. Prod. Res.* **49**(23), 6963–6980 (2011). <https://doi.org/10.1080/00207543.2010.526016>
  15. Vieira, S., Lopes, R.B.: Improving production systems with lean: a case study in a medium-sized manufacturer. *Int. J. Ind. Syst. Eng.* **33**(2), 162–180 (2019). <https://doi.org/10.1504/IJISE.2019.10024260>
  16. Laborie, P., Rogerie, J., Shaw, P., Vilím, P.: IBM ILOG CP optimizer for scheduling: 20+ years of scheduling with constraints at IBM/ILOG. *Constraints* **23**(2), 210–250 (2018). <https://doi.org/10.1007/s10601-018-9281-x>
  17. Brandimarte, P.: Routing and scheduling in a flexible job shop by tabu search. *Ann. Oper. Res.* **41**(3), 157–183 (1993). <https://doi.org/10.1007/BF02023073>
  18. Kacem, I., Hammadi, S., Borne, P.: Pareto-optimality approach for flexible job-shop scheduling problems: hybridization of evolutionary algorithms and fuzzy logic. *Math. Comput. Simul.* **60**(3–5), 245–276 (2002). [https://doi.org/10.1016/S0378-4754\(02\)00019-8](https://doi.org/10.1016/S0378-4754(02)00019-8)
  19. Ferreira, J.: Otimização e apoio à decisão no escalonamento da produção de cápsulas de rosca vinícolas em ambiente flexible job shop. MSc thesis. University of Aveiro, Aveiro, Portugal (2020). <http://hdl.handle.net/10773/31720>

# Bimaterial Three-Dimensional Printing Using Digital Light Processing Projectors



Daniel Bandeira and Marta Pascoal

**Abstract** The bimaterial three-dimensional (3D) printing problem focuses on printing 3D parts with two materials, one that forms an inner structure and another one that coats it. The inner structure forms a shaded area that compromises traditional 3D printing techniques. This work approaches such a problem, assuming that Digital Light Processing (DLP) projectors can be installed on the printer walls. We consider that the image emitted by the DLP projector is used to cure each layer and that several of these devices are installed with the purpose of reaching shaded ones. The goal is bifold: to maximize the extension of the cured object and to minimize its distortion, resulting from possible over-exposed outer areas. The problem is divided into two subproblems: the problem of locating projectors around the printer to maximize the polymerized part of a given object, and the problem of assigning the projectors selected in the previous step to the areas of the resin to be cured, while minimizing the areas that result from the distortion of the printing. Integer linear formulations and resolution methods are presented for both. Algorithms are tested and the results are discussed. For the considered case study, a solution was found with 18.08% of voxels left uncured when using 1 DLP projector. The same percentage decreases to 0.90, or 0.05%, when using 3, or 6, DPL projectors. According to the results for the second subproblem, there should always be a projector on the top of the printer, in addition to others on the side walls. In terms of distortion, the results show that it is convenient to have at least 4 DLP projectors and, in such case, the inner area left uncured and the outer area over-exposed are both below 2.5%. When using 6 emitters, the first problem required 4 s of run time, and the second required up to 3 min.

---

D. Bandeira · M. Pascoal (✉)

University of Coimbra, CMUC, Department of Mathematics, 3001-501 Coimbra, Portugal

e-mail: [marta@mat.uc.pt](mailto:marta@mat.uc.pt)

M. Pascoal

Institute for Systems Engineering and Computers, Coimbra, Portugal

Dipartimento di Elettronica, Informazione e Bioingegneria, Politecnico di Milano,

Piazza Leonardo da Vinci, 32, 20133 Milan, Italy

**Keywords** Combinatorial optimization · Maximum coverage problem applications · Biobjective optimization · 3D printing

## 1 Introduction

Three-dimensional printing, or 3D printing, is an additive process for rapid free form manufacturing, where the final object is created by the addition of thin layers of material. Each layer is a cross-section of the object to be constructed, and the printer draws each layer as a 2D printing. Details on these processes can be found in [3, 8]. The technology and materials required for this type of processes started being developed around 1980. Several improvements have been made to 3D printing methods since their origin. Nowadays those printers are fairly affordable, which has made the process popular due to its ability to produce objects (often called components or parts) quickly and at a low cost [10].

One of the technologies for 3D printing is stereolithography (SL). The object to produce is divided into layers that are successively added as a liquid polymer (a resin) along the printing process. Each 2D layer is partitioned into uniform squares, called voxels, to specify the part to construct. The zone of the layer reached by an ultraviolet (UV) laser light is then polymerized (or cured). Then, the platform that supports the model moves to prepare for printing the next layer. The voxels correspond to a discretization of the part. Naturally, the more and the smaller the voxels, the more accurate the printing. When the printing is over, the part is clean from remaining liquid resin and goes through a post-cure process in a UV oven for a few minutes, to fully cure the resin.

Traditional 3D SL processes use a single fixed emitter of laser light, or a system of galvanometer mirror scanners to project the laser beam in the wished directions. As an alternative, Digital Light Processing (DLP) uses UV light emitted from a video projector, known as a DLP projector. To simplify, a light source used to cure the resin will sometimes be called an emitter. The projector flashes a single image of each layer across the printing platform at once to cure the liquid resin. Because the projector is a digital screen, the image of each layer is formed by square pixels, which results in a layer of voxels, rather than in sequentially drawing the image as described above. In general, DLP printing is faster than SL and uses fewer moving parts than SL with galvanometer mirror scanners. The accuracy of the produced model depends on the resolution of the projector. Details on additive manufacturing and on different forms of SL can be found in [14].

In many cases, the use of a single material in 3D printing has limitations in terms of the mechanical properties of the part [11, 16]. Composite structures involving more than one material have been developed to overcome such limitations, exploiting aspects such as the materials or the patterns to enhance the overall mechanical properties, like stiffness, toughness, ductility, and impact resistance. Multimaterial components have a wide range of applications to custom orthotics, intelligent components, complex or fragile parts where over-injection or others options are not feasible

or are not economically sustainable [15, 18]. A review on this topic can be found in [6]. Some of the approaches to multimaterial 3D printing explore the possibility of printing one layer at a time, like in the traditional 3D printing, but splitting the two (or more materials) [9, 17]. Nevertheless, to our knowledge, the problem has merited little attention.

The current work follows a different approach from the previous, by considering a process analogous to SL for printing an object in which the polymer coats a constructed 3D grid structure made of a different material, like metal or ceramic. This procedure can be extended to more than two materials. The existence of a grid that supports the polymer raises some difficulties, as it may block the light from the emitter, thus creating shaded and uncured areas. We call this problem the bimaterial 3D printing problem. Its goal is to study how to overcome this issue and print the part with as little distortion as possible. A previous work addressed this question when using the traditional SL methodology, that is, using galvanometer mirror scanner systems for curing the resin [1]. The light beam emitted by such devices has a circular shape, which may become elliptical when it intersects a plane. The galvanometer mirrors can be redirected to different regions by changing its angle of incidence in the printing plane. The problem was approached in two steps. The first consists of minimizing the number of emitters to install in the printer in a way to reach all the voxels of the part. The second considers independently the printing of the different layers and it consists of assigning the emitters to the voxels with two goals: maximizing the angle of incidence of the beam onto the printing plane, in order to minimize the distortion of the part, while minimizing the number of emitters involved in the printing.

In the current work we consider devices with different characteristics as the light sources, and this requires different modelling and resolution approaches. The use of DLP projectors projects an image formed by several pixels on the printing platform. In this case the pixels have a square shape and their distorted versions are parallelograms, which can be corrected to rectangles if using a special lens. In addition, the DLP projectors have fixed positions along the printing process and the decision to be made for each layer is which pixels to switch on. Assuming that both the voxels to cure and the possible locations for the DLP projectors are known, the problem is formulated in two phases:

- Emitters location problem (ELP): aiming at finding the emitters' positions, with the goal of maximizing the number of voxels that can be cured.
- Emitters assignment problem (EAP): aiming at assigning each voxel to cure with a DLP projector installed according to the solution of the ELP, to obtain the best possible print.

Ideally, maximizing the number of voxels reached corresponds to reaching all of them. However, the inner unreached voxels can be fixed with this post-cure procedure, therefore, other solutions may still be acceptable. The ELP is modelled as a maximum coverage problem and greedy heuristics are described to solve it. The EAP is modelled as an integer linear program (ILP) with two objective functions, which reflect two areas related with the distortion of the light emitted by the DLP

projector when hitting the component voxels. The distortion of the printed object is simpler to assess than with the technique used in [1], therefore this is integrated with the method proposed for this phase of the problem. A greedy exact algorithm based on a weighted-sum method is presented to find the supported non-dominated points for the problem. The methodology developed for the two phases is applied to a case study and the behavior of the new methods is assessed in terms of the output solutions and of the run times.

The rest of the paper is organized as follows. The ELP and the EAP are addressed in Sects. 2 and 3. The problems are formulated and algorithms for solving them are presented. Section 4 is dedicated to computational results of the methods for a case study. Concluding remarks are drawn in the last section.

## 2 Emitters Location Problem

Let  $M = \{1, \dots, m\}$  be the set of voxels of the part and  $N = \{1, \dots, n\}$  be the set of possible positions for the emitters. The emitter at position  $j$ , or simply emitter  $j$ , is said to cover the voxel  $i$  if it can reach it with the laser light, for  $i \in M, j \in N$ . Let the emitters coverage matrix be  $A = [a_{ij}]_{i=1, \dots, m; j=1, \dots, n}$ , such that

$$a_{ij} = \begin{cases} 1 & \text{if the emitter } j \text{ can reach the voxel } i \\ 0 & \text{otherwise} \end{cases}, \quad i \in M, \quad j \in N.$$

We assume this matrix is given, as it can be obtained by the ray tracing algorithm [12].

Let  $x_j$  and  $y_i$  be binary decision variables such that  $x_j = 1$  if and only if the emitter  $j$  is installed, and  $y_i = 1$  if and only if the voxel  $i$  is reached,  $i \in M, j \in N$ . Then, the total number of reached voxels, to be maximized, is given by

$$\sum_{i=1}^m y_i.$$

Let  $b$  be the number of emitters that can be used for printing the object. Then,

$$x_1 + \dots + x_n \leq b \tag{1}$$

should hold. Finally, it is necessary to link the variables  $x_j$  and  $y_i, i \in M, j \in N$ . To do so, a given voxel will be considered as reached if it can be reached by at least one emitter in a selected position, through the constraints

$$\sum_{j=1}^n a_{ij} x_j \geq y_i, \quad i \in M,$$

or, in matrix notation,

$$Ax \geq y. \quad (2)$$

Then, the ELP problem can be formulated as

$$\begin{aligned} \max \quad & \sum_{i=1}^m y_i \\ \text{subject to} \quad & (1), (2) \\ & x \in \{0, 1\}^n, \ y \in \{0, 1\}^m \end{aligned} \quad (\text{P1})$$

The optimal solution of (P1) provides the positions where the light emitters should be installed (i.e., the indexes  $j$  such that  $x_j = 1$ ), as well as the positions of voxels that are reached (i.e., the indexes  $i$  for which  $y_i = 1$ ). This is a coverage problem, with the goal of covering the maximum number of voxels with a bounded number of projectors. This problem was shown to be NP-hard, thus exact algorithms may not be able to solve large instances. Instead, two greedy heuristic methods are outlined as alternatives for finding solutions to problem (P1). The first is based on the information obtained from the columns of the emitters coverage matrix, the possible positions for the emitters. The second is based on the information from both the rows and the columns of the emitters coverage matrix, the voxels to cure and the possible emitters.

Let  $N_i = \{j \in N : a_{ij} = 1\}$  be the set of emitters that can reach the voxel  $i \in M$ , and  $M_j = \{i \in M : a_{ij} = 1\}$  be the set of the voxels that can be reached by the emitter  $j \in N$ . The first heuristic is based on selecting first the emitters that cover the most uncovered voxels. The algorithm uses two auxiliary variables:  $S$ , a set that stores the indexes of the selected emitters; and  $R$ , a set that stores the indexes of the uncovered voxels. A voxel is removed from  $R$  whenever an emitter that covers it is selected. The process stops when  $b$  emitters are selected or no more voxels can be reached. The set  $R$  indicates the voxels that are left unreached, whereas the set  $S$  provides a solution. Algorithm 1 outlines this method.

---

**Algorithm 1:** Greedy heuristic based on the emitters

---

```

1  $S \leftarrow \emptyset$ 
2  $R \leftarrow M$ 
3  $k \leftarrow 0$ 
4  $continue \leftarrow \text{true}$ 
5 while  $k < b$  and  $continue$  do
6    $j^* \leftarrow \text{argmax}\{|M_j \cap R| : j \in N \setminus S\}$ 
7   if  $j^*$  is undefined then
8      $continue \leftarrow \text{false}$ 
9   else
10     $S \leftarrow S \cup \{j^*\}$ 
11     $R \leftarrow R \setminus M_{j^*}$ 
12     $k \leftarrow k + 1$ 
13 if  $R \neq \emptyset$  then The emitters in  $S$  do not reach all the voxels

```

---

Algorithm 1 corresponds to a classical greedy method for the general version of the maximum coverage problem. Nemhauser et al. [13] proved this heuristic to have an approximation ratio of  $(1 - \frac{1}{e})$ .

In the second heuristic, the voxel that is covered by fewer emitters is identified, since it is the hardest to cover. Then, like in Algorithm 1, among all the emitters that cover the chosen voxel, the one that covers more voxels is preferred. This method uses the same auxiliary sets defined above and is outlined in Algorithm 2.

---

**Algorithm 2:** Greedy heuristic based on the voxels and the emitters

---

```

1  $S \leftarrow \emptyset$ 
2  $R \leftarrow M$ 
3  $k \leftarrow 0$ 
4  $continue \leftarrow \text{true}$ 
5 while  $k < b$  and  $continue$  do
6    $i^* \leftarrow \operatorname{argmin} \left\{ \sum_{j \in N_i} a_{ij} : i \in R \right\}$ 
7    $j^* \leftarrow \operatorname{argmax} \{ |M_j \cap R| : j \in N_{i^*} \setminus S \}$ 
8   if  $j^*$  is undefined then
9      $continue \leftarrow \text{false}$ 
10  else
11     $S \leftarrow S \cup \{j^*\}$ 
12     $R \leftarrow R \setminus M_{j^*}$ 
13     $k \leftarrow k + 1$ 
14 if  $R \neq \emptyset$  then The emitters in  $S$  do not reach all the voxels

```

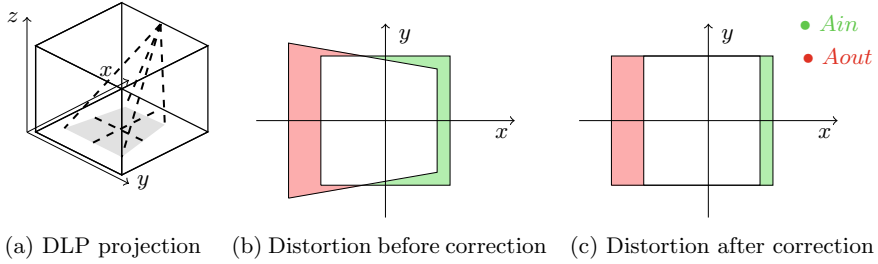
---

Both Algorithms 1 and 2 have time complexity of  $O(b(n + m))$ .

### 3 Emitters Assignment Problem

The next step for printing the object is to assign emitters and voxels, which is the goal of the EAP. The emitters are at fixed positions while printing the object. Furthermore, the EAP focuses on a layer at a time, as these are independent printings. Let us consider that for any layer the emitters coverage matrix,  $A$ , is restricted to the  $b$  positions selected for the emitters. This is done without loss of generality because the number of emitters used for printing depends on the number of emitters installed, rather than on the maximum number of emitters allowed. Let us also assume that the  $c$  voxels that can be reached on that layer. That is, in this section  $A = [a_{ij}]_{i=1,\dots,c; j=1,\dots,b}$ , with  $a_{ij} = 1$  if and only if the emitter  $j$  can reach the voxel  $i$ .

The assignment of voxels and emitters is related with the quality of the part to print, which depends on the distortion of the light emitted by the DLP projector when it hits the voxels. Each DLP projector projects an image formed by a matrix of pixels that can be either on or off—Fig. 1a. If the projector is not perpendicular to



**Fig. 1** Projection of a DLP over the printing platform

the printing platform, the light emitted by each pixel is distorted as shown in Fig. 1b. However, special lenses can correct the produced trapezoid to the shape of a rectangle as shown in Fig. 1c.

The green and red areas in Fig. 1c are called inner ( $A_{in}$ ) and outer ( $A_{out}$ ) areas, respectively, and reflect the effect of the light distortion when curing a voxel of resin. The inner area corresponds to the part of the voxel that is uncured because it is not reached by UV light, whereas the outer area corresponds to the resin that is cured beyond the voxel due to the light distortion.

The total value of the inner and outer areas is used to estimate the overall quality of the printed part; therefore, these are computed in the following. In general, these calculations are over-estimations of the actual values, because the unreached area of a voxel may be reached by the distorted light coming from another pixel, and because the outer area is only relevant for the voxels in the border of the part. Let us consider the parameters below, shown in Fig. 2: the height of the DLP projector ( $h$ ); the projector's incidence angle ( $\alpha$ ); the angle of the projector in the direction of the  $yy$  axis ( $\beta$ ). Point  $O$  is the center of the projection and  $P$  marks the projector's position. While the projector is at a fixed position, it is considered that it can tilt only in the direction of the  $yy$  axis.

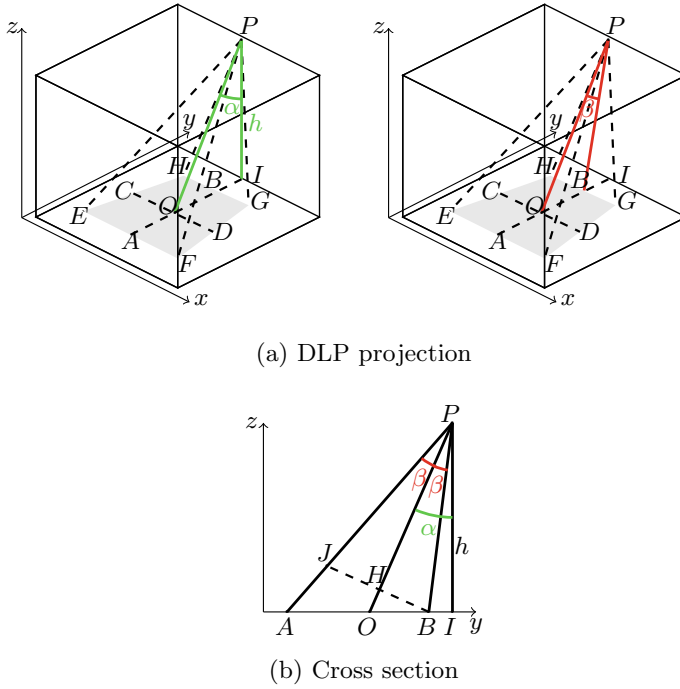
The DLP projector image can be corrected along the  $xx$  axis, and then have the overall shape of a rectangle. If the projection is perpendicular, the pixels produced are uniform squares. However, in the general case there are differences in the rectangles produced by the different pixels of the projector. The side of the squares obtained perpendicularly will be used as the measure unit hereafter.

Figure 2b shows the cut of the plane  $IPA$  in Fig. 2a. The segment  $\overline{JB}$  corresponds to the case of a projection perpendicular to the printing surface, where pixels produce squares of the same size in this segment. Then,

$$\overline{OB} = \overline{OI} - \overline{BI} = h \tan \alpha - h \tan(\alpha - \beta)$$

and  $\hat{ABJ} = \alpha$ , therefore,

$$\overline{JB} = 2h \cos \alpha (\tan \alpha - \tan(\alpha - \beta)).$$



**Fig. 2** Projection and its cross section

If the row projected in the direction of the  $yy$  axis has  $p$  pixels, then each pixel in  $\overline{JB}$  is projected as a square with side length given by  $\frac{\overline{JB}}{p} = \frac{2h \cos \alpha}{p} (\tan \alpha - \tan(\alpha - \beta))$ . Moreover,  $\tan \beta = \frac{\overline{HB}}{\overline{HP}}$ , thus,

$$\overline{HP} = \frac{h \cos \alpha (\tan \alpha - \tan(\alpha - \beta))}{\tan \beta}.$$

Therefore, the length of the side of the rectangle for the projection of pixel  $i$  starting at  $B$  is

$$l_i = h \tan \left( \alpha - \arctan \left( \left( 1 - \frac{2i}{p} \right) \tan \beta \right) \right) - h \tan \left( \alpha - \arctan \left( \left( 1 - \frac{2(i-1)}{p} \right) \tan \beta \right) \right)$$

for  $i = 1, \dots, \frac{p}{2}$ , and,

$$l_i = h \tan \left( \alpha + \arctan \left( \left( \frac{2i}{p} - 1 \right) \tan \beta \right) \right) - h \tan \left( \alpha + \arctan \left( \left( \frac{2(i-1)}{p} - 1 \right) \tan \beta \right) \right)$$

for  $i = \frac{p}{2} + 1, \dots, p$ . These formulas hold for perpendicular projections, that is, when the projector is placed above and at the center of the projection plane, the

first one for pixels to its left (with positive angles) and the second for pixels on the right (with negative angles). Therefore, knowing the coordinates of the projector, its angles  $\alpha, \beta$  can be determined. To do this we consider  $y_1$ , the ordinate of point  $O$  in Fig. 2a. Recall that  $O$  is the point at which the projector points. If  $y_2$  is the ordinate of point  $P$ , then, on the one hand, the sum of the sides of the pixels along the  $yy$  axis must be equal to  $p$ , and, on the other, the sum of the sides of the first  $\frac{p}{2}$  pixels along the same axis must be  $y_1 - y_2$ . Then,  $y_1$  and  $\beta$  can be found by solving

$$\begin{cases} l_1 + \dots + l_p = p \\ l_1 + \dots + l_{\frac{p}{2}} = y_1 - y_2 \end{cases}$$

Considering that point  $P$  has coordinates  $P = (p_1, p_2, p_3)$ , the value of  $\alpha$  is then given by

$$\alpha = \arctan\left(\frac{p_2 - y_1}{p_3}\right).$$

Given the angles  $\alpha, \beta$ , the values of the sides of all pixels can be calculated.

In general, several projectors may cure the resin at a given voxel, so the choice of one of them depends on the areas  $A_{in}$  and  $A_{out}$ . Let  $A_{in_{ij}}$  ( $A_{out_{ij}}$ ) denote the inner (outer) area produced by the emitter  $j$  over the voxel  $i$ . These values are calculated choosing the nearest pixel of emitter  $j$  to voxel  $i$ , because this is the one that produces the least distortion,  $i = 1, \dots, c$ ,  $j = 1, \dots, b$ . The goal of the EAP is to assign the projectors to activate to the voxels of the part, in a way that minimizes the distortion. To formulate the problem, let  $y_{ij}$  be binary decision variables with value 1 if and only if emitter  $j$  is activated to reach voxel  $i$ ,  $i = 1, \dots, c$ ,  $j = 1, \dots, b$ .

The assignment between emitters and voxels must take two aspects into account: the uniqueness of the solution and its viability. The first is ensured by

$$\sum_{j=1}^b y_{ij} = 1, \quad i = 1, \dots, c, \quad (3)$$

while the second depends on the conditions

$$\sum_{j=1}^b a_{ij} y_{ij} = 1, \quad i = 1, \dots, c. \quad (4)$$

The two objective functions represent the areas to be minimized,

$$z_1(y) = \sum_{i=1}^c \sum_{j=1}^b A_{in_{ij}} y_{ij} \quad \text{and} \quad z_2(y) = \sum_{i=1}^c \sum_{j=1}^b A_{out_{ij}} y_{ij}.$$

Thus, the following formulation is obtained for the EAP

$$\begin{aligned}
& \min && z_1(y) \\
& \min && z_2(y) \\
& \text{subject to} && (3), (4) \\
& && y \in \{0, 1\}^{c \times b}
\end{aligned} \tag{P2}$$

It is worth noting that both constraints (3) and (4) are needed to ensure that a single emitter reaches each voxel and also that the voxel can be reached by the emitter. If the constraints (3) are omitted, we may have an optimal solution  $y$  such that  $y_{ij} = y_{ij'} = 1$  for a given voxel  $i$  and two emitters  $j, j'$  such that  $a_{ij} \neq a_{ij'}$  and  $Ain_{ij} = Aout_{ij} = 0$  or  $Ain_{ij'} = Aout_{ij'} = 0$ . As an alternative, a post-processing routine can be applied to correct such solutions, in case the constraints (3) are omitted.

The two objective functions are expected to be conflicting, thus, in general, there is no feasible solution that simultaneously optimizes both. As a consequence, when dealing with two objectives we seek to determine Pareto optimal, or efficient, solutions. A feasible solution of (P2),  $y$ , is said to be efficient if there is no other feasible solution,  $x$ , that dominates it, that is, such that

$$z_1(x) \leq z_1(y), \quad z_2(x) \leq z_2(y), \quad (z_1(x), z_2(x)) \neq (z_1(y), z_2(y)).$$

The image of an efficient solution is a point in the criteria space, called a non-dominated point [7]. The efficient solutions of a biobjective problem can be classified into supported and non-supported. The first are optimal solutions of a weighted-sum problem that minimizes the function  $\lambda z_1 + (1 - \lambda)z_2$ , with  $0 < \lambda < 1$ . The remaining efficient solutions are non-supported, or unsupported, and the same holds for non-dominated points.

Three approaches to multiobjective optimization are usually considered: a priori methods, when the decision maker expresses the preferences before calculating the solutions, for example by means of a single objective utility function; interactive methods, with a progressive articulation of the preferences of the decision maker as the solutions are presented, and a posteriori methods, which calculate all solutions, without the need for a definition of preferences. The relative importance of the two objective functions in the EAP is unclear, so an a posteriori method is the most appropriate. Several methods are known to find efficient solutions or non-dominated points of biobjective ILPs, for instance, [2, 4]. For the bimaterial 3D printing application it is not required to know all the efficient solutions, therefore we adapt the noninferior set estimation (NISE) method for approximating the set of supported non-dominated points [5].

Most multiobjective optimization methods solve a sequence of single objective problems related with the original. In the NISE method those problems are subject to the same constraints and the new objective function is a weighted-sum of the initial objective functions. The weights depend on the adjacent computed images, and are chosen in a way that allows to explore the images space. The first pair of solutions is formed by the lexicographic optimal solutions for the two objective functions,  $(z_1, z_2), (z_2, z_1)$ . Afterwards, for any pair of adjacent points,  $A = (z_1(y_A), z_2(y_A))$ ,  $B = (z_1(y_B), z_2(y_B))$ , the weighted-sum problem

$$\begin{aligned}
& \min \quad z(y) = \lambda_1 z_1(y) + \lambda_2 z_2(y) \\
& \text{subject to (3), (4)} \\
& y \in \{0, 1\}^{c \times b}
\end{aligned} \tag{P3}$$

is solved, with the weights  $\lambda_1 = z_2(y_A) - z_2(y_B)$ ,  $\lambda_2 = z_1(y_A) - z_1(y_B)$ , which can be normalized to the interval  $[0, 1]$ . If no new solution is obtained when solving this problem, no new pair of solutions is considered. The algorithm halts when all the pairs of solutions have been scanned. If not all supported non-dominated points need to be known, the algorithm can be relaxed by discarding new solutions whenever they are closer than a small  $\epsilon > 0$ . The method is outlined in Algorithm 3. The variable  $P$  stores the supported non-dominated points found by the algorithm, while  $Y$  stores pairs of supported non-dominated points yet to scan.

---

**Algorithm 3:** Weighted-sum method to approximate the set of supported non-dominated points

---

```

1  $y_A \leftarrow$  lexicographic optimal solution of (P2) with respect to  $(z_1, z_2)$ 
2  $y_B \leftarrow$  lexicographic optimal solution of (P2) with respect to  $(z_2, z_1)$ 
3  $P \leftarrow \{A, B\}; Y \leftarrow \{(A, B)\}$ 
4 while  $Y \neq \emptyset$  do
5    $(A, B) \leftarrow$  element of  $Y$ 
6    $Y \leftarrow Y - \{(A, B)\}$ 
7    $\lambda_1 \leftarrow z_2(y_A) - z_2(y_B)$ 
8    $\lambda_2 \leftarrow z_1(y_A) - z_1(y_B)$ 
9    $y_C \leftarrow$  optimal solution of (P3)
10  if  $|z(y_C) - z(y_A)| \geq \epsilon$  and  $|z(y_C) - z(y_B)| \geq \epsilon$  then
11     $P \leftarrow P \cup \{C\}; Y \leftarrow Y \cup \{(A, C), (C, B)\}$ 

```

---

A closed form for the solution of problem (P3), used on line 9 of Algorithm 3, can be derived. Since the values  $Ain_{ij}$ ,  $Aout_{ij}$  are known,  $z$  can be written as

$$z(y) = \sum_{i=1}^c \sum_{j=1}^b Atot_{ij} y_{ij},$$

where  $Atot_{ij} = \lambda_1 Ain_{ij} + \lambda_2 Aout_{ij}$ ,  $i = 1, \dots, c$ ,  $j = 1, \dots, b$ . Constraints (3) and (4) require that only one of the variables  $y_{ij}$  is 1, while the remaining are 0, for each voxel  $i = 1, \dots, c$ . Then, the objective function is minimum when

$$y_{ij} = \begin{cases} 1 & \text{if } j = \operatorname{argmin}\{Atot_{ij} : a_{ij} = 1\} \\ 0 & \text{otherwise} \end{cases}, \quad i = 1, \dots, c, \quad j = 1, \dots, b.$$

This method can also be adapted to determine the lexicographic optimal solutions using as argument of  $\operatorname{argmin}$  the preferred function. In case of a tie for some emitters, the one with the lowest value for the second function is chosen.

## 4 Computational Experiments

The methods proposed for the ELP and the EAP are now tested for a case study. The approaches are compared and assessed in terms of the output solutions and run times. However, the latter is not a major aspect for the bimaterial 3D printing problem, given that both ELP and EAP are solved before printing.

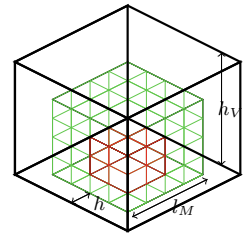
The case study is retrieved from that considered in [1] and it consists of a cube with a hole at the bottom. This inner hole is also a cube, embedded in the first, as shown in Fig. 3. The grid in the image (in green and red) represents the existing metal grid, to be coated with resin. The hole is at the base of the cube (in red). The thickness of the metal grid is considered to be equal to the thickness of the polymer layers, which is assumed to be a generic length unit of 1. This value is also the width of the voxels and it corresponds to 0.2 mm. The variable parameters for printing, expressed in terms of generic length units, are:

- The length of each segment of the metal grid,  $l_M$ .
- The thickness of the polymer added on each side of the metal grid,  $l_P$ .
- The number of divisions of the metal grid, assumed to be uniform,  $n_M$ .
- The distance between the cube and the side walls of the printer,  $h$ .
- The height of the printing area,  $h_V$ .
- The length of the side of the hole,  $n_b$ .

For all faces, the length of the edge of the final cube to be printed corresponds to the addition of polymer on each side to the length of the edges of the cube of metal, i.e.  $n_V = 2l_P + l_M$ . Thus, the object has  $2l_P + l_M$  layers. The upper layers are excluded from the tests, because the emitter on the top of the printer can be used directly with a traditional 3D printing process. Moreover, the cube to print is assumed to be centered on the printing platform. In addition, the printer dimensions are fixed to  $h_V = 1250$  and  $h = 475$ , and the produced cube has side  $n_V = 300$ . We also consider the fixed value  $l_P = 1$ . Each layer has  $n_V \times n_V$  voxels. The remaining characteristics of the solved instances are shown in Table 1.

Algorithms 1, 2, and 3 were implemented in MATLAB R2016b. The tests run on an Intel i7-6700 Quadcore of 3.4 GHz, with 8 Mb of cache and 16 Gb of RAM. The run times are average values obtained for 30 repetitions, for accuracy and trying to minimize the effect of possible oscillations of the CPU.

**Fig. 3** Printing area and object to be printed ([1])



**Table 1** Test parameters

| Test  | T1 | T2 | T3 | T4 |
|-------|----|----|----|----|
| $n_M$ | 10 | 10 | 12 | 12 |
| $n_b$ | 2  | 4  | 2  | 4  |

4.1 Emitters Location Problem

For solving the ELP, 107 possible emitter locations were considered. These positions are evenly distributed on the top and the side walls of the printer. A minimal height of 500 was imposed along the printer side walls, with the goal of bounding the distortion of the light on the printing layer.

The average results for  $b = 1, \dots, 6$  emitters are summarized for Algorithms 1 and 2 in Table 2, or in Figs. 4 and 5 in Appendix 6. The Algorithms 1 and 2 produced very similar solutions in terms of the percentage of voxels that cannot be cured with the obtained solutions ( $\mu$ ). The average percentages are always below 20% and, as expected, decrease significantly when the number of installed DLP projectors increases, especially for 3 or more projectors, but even more for 5 projectors. The number of uncured voxels is also higher when  $n_M$  and  $n_b$  are bigger, although this is particularly clear with a single projector. Algorithm 1 outperformed Algorithm 2 for these instances, so, overall, it seems a better choice than the latter.

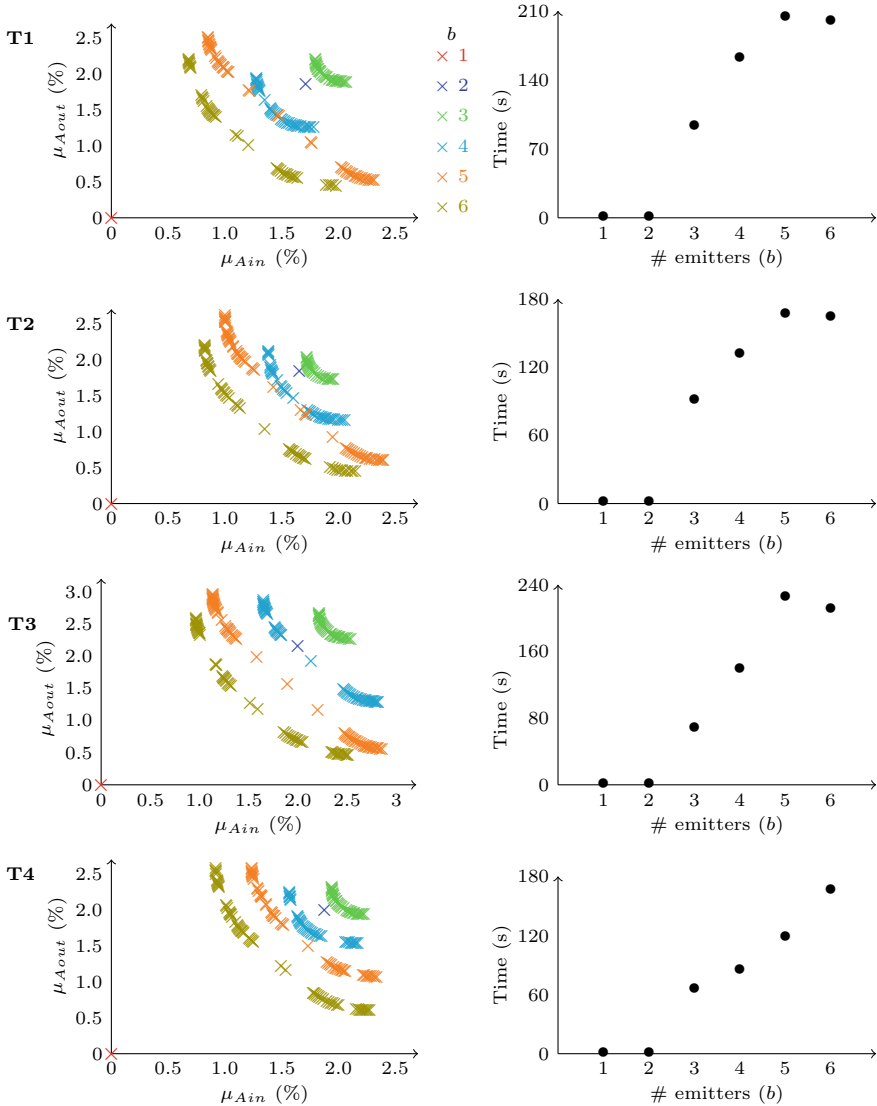
**Table 2** Average results for the ELP

| Percentage of voxels left to be cured |             |      |      |      |      |      |             |      |      |      |      |      |
|---------------------------------------|-------------|------|------|------|------|------|-------------|------|------|------|------|------|
| # emitters                            | Algorithm 1 |      |      |      |      |      | Algorithm 2 |      |      |      |      |      |
| (b)                                   | 1           | 2    | 3    | 4    | 5    | 6    | 1           | 2    | 3    | 4    | 5    | 6    |
| T1                                    | 13.91       | 3.97 | 0.42 | 0.15 | 0.03 | 0.02 | 14.77       | 5.26 | 0.72 | 0.15 | 0.09 | 0.03 |
| T2                                    | 15.54       | 4.46 | 0.57 | 0.19 | 0.03 | 0.03 | 16.63       | 2.69 | 0.63 | 0.09 | 0.05 | 0.03 |
| T3                                    | 16.73       | 5.16 | 0.83 | 0.32 | 0.05 | 0.04 | 16.83       | 2.52 | 0.63 | 0.07 | 0.05 | 0.04 |
| T4                                    | 18.08       | 5.63 | 0.90 | 0.34 | 0.09 | 0.05 | 19.43       | 4.07 | 1.72 | 0.81 | 0.16 | 0.06 |
| Run times (s)                         |             |      |      |      |      |      |             |      |      |      |      |      |
| # emitters                            | Algorithm 1 |      |      |      |      |      | Algorithm 2 |      |      |      |      |      |
| (b)                                   | 1           | 2    | 3    | 4    | 5    | 6    | 1           | 2    | 3    | 4    | 5    | 6    |
| T1                                    | 1.86        | 2.39 | 2.71 | 2.57 | 2.08 | 2.15 | 2.60        | 3.55 | 3.65 | 3.64 | 4.05 | 4.01 |
| T2                                    | 2.35        | 3.24 | 2.88 | 3.09 | 2.94 | 2.93 | 2.98        | 3.14 | 4.04 | 3.69 | 4.35 | 4.53 |
| T3                                    | 2.12        | 2.20 | 2.63 | 2.34 | 2.17 | 2.78 | 3.09        | 4.08 | 3.85 | 3.81 | 4.29 | 3.73 |
| T4                                    | 1.79        | 3.33 | 2.81 | 3.00 | 2.34 | 2.91 | 3.34        | 3.81 | 4.18 | 4.34 | 4.44 | 4.80 |

## 4.2 Emitters Assignment Problem

After the location of the emitters is known, the EAP can be solved. Given the findings of Sect. 4.1, Algorithm 3 was applied to the solutions of Algorithm 1.

Figure 4 shows the average percentages for the inner and outer voxel areas,  $\mu_{Ain}$ ,  $\mu_{Aout}$ , for the supported non-dominated points found by Algorithm 3 when



**Fig. 4** EAP: Average  $Ain$ ,  $Aout$  for supported points and run times

constructing all layers. The EAP has a single solution when using 1 or 2 emitters. Furthermore, with  $b = 1$  the value of both areas is 0, because in this case the emitter is on the top wall and, thus, hits the printing layer perpendicularly. Moreover, when  $b = 2$  the solution is to install one emitter on the top wall, to prevent image distortion, and select the second in order to reach the voxels that cannot be cured with the first. For this reason, the solution is unique when  $b = 2$ . For the remaining cases, more non-dominated points are found and the values of  $A_{in}$ ,  $A_{out}$  decrease as  $b$  increases. There is also a small increase of  $\mu_{A_{in}}$  for tests T3 and T4, when the initial grid is denser, which may be explained by the fact that more voxels may be uncured in these cases. Still concerning the overall values of  $A_{in}$ ,  $A_{out}$  found for the EAP, it seems to be advantageous to use the highest possible number of DLP projectors (at least 4).

In general, Algorithm 3 is slower when  $b$  increases, but the average run times never exceed 4 min. As stressed before, bimaterial 3D printing is not a real time application, so the run times of the method are suitable for this purpose.

## 5 Conclusions

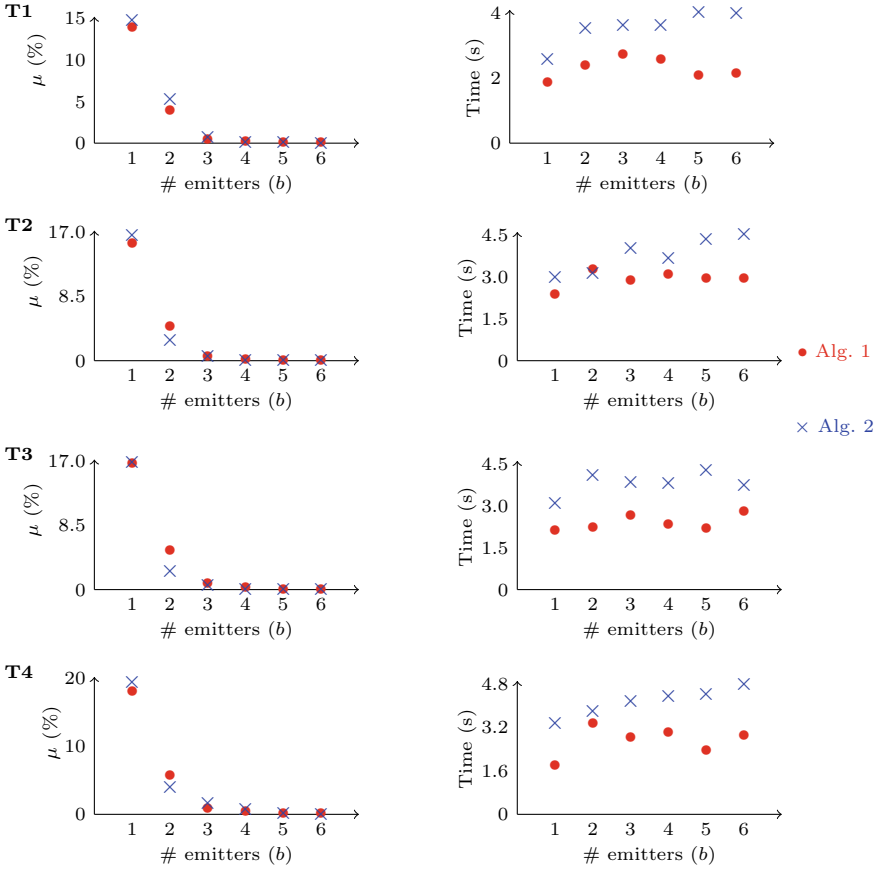
We addressed the problem of installing DLP projectors for printing 3D bimaterial parts. Two steps are considered: locating the emitters and assigning them to the voxels of the part to construct. The first is modelled as a coverage problem and is solved by greedy heuristics. The second is formulated as an ILP that minimizes the uncured area and the area of resin that is cured beyond the pre-defined voxels, whose supported non-dominated points are found by a weighted-sum method.

Solutions for the ELP were found by heuristics in few seconds and have few unreached voxels. However, that number decreases significantly when using 3 or more emitters. Solutions for the EAP were also found in up to 4min, which is reasonable for planning problems. The values for the areas of the uncured voxels and the voxels that should not have been cured depend on the number of emitters and almost always decrease when more emitters are used. Nevertheless, their values do not exceed 3%, even if these are upper bounds for the actual values. A decision maker will have to choose a single solution among the several found for the EAP. Future research may consider the emitters to be in fixed positions but able to rotate in other directions besides the vertical, or consider that those positions change between different layers.

**Acknowledgements** Developed within project PT2020-POCI-SII & DT 17963: NEXT.Parts, Next-Generation of Advanced Hybrid Parts, through the COMPETE 2020-POCI. Partially supported by the Portuguese Foundation for Science and Technology (FCT) under project grants UID/MAT/00324/2020 and UID/MULTI/00308/2020.

## 6 Appendix: Computational Results

See Fig. 5.



**Fig. 5** ELP: solutions (uncured voxels) and run times of the Algorithms 1 and 2

## References

1. Bandeira, D., Pascoal, M., Santos, B.: Bimaterial 3D printing using sterolithography with galvanometer scanners. *Optim. Eng.* **21**, 49–72 (2020)
2. Boland, N., Charkhgard, H., Savelsbergh, M.: A criterion space search algorithm for biobjective integer programming: the balanced box method. *INFORMS J. Comput.* **2**, 558–584 (2015)
3. Burns, M.: *Automated Fabrication: Improving Productivity in Manufacturing*. Prentice-Hall, Inc. (1993)
4. Clímaco, J., Pascoal, M.: An approach to determine unsupported non-dominated solutions in bicriteria integer linear programs. *INFOR: Inf. Syst. Oper. Res.* **54**, 317–343 (2016)
5. Cohon, J.L.: *Multiobjective Programming and Planning*. Academic Press, New York (1978)
6. Cooperstein, I., Sachyani-Keneth, E., Shukrun-Farrell, E., Rosental, T., Wang, X., Kamyshny, A., Magdassi, S.: Hybrid materials for functional 3D printing. *Adv. Mater. Interfaces* **5**, 1800996 (2018)
7. Ehrgott, M.: *Multicriteria Optimization*. Springer Science & Business Media, Berlin, Heidelberg (2006)
8. Gibson, I., Rosen, D., Stucker, B.: *Additive Manufacturing Technologies*. Springer, New York (2015)
9. Kang, W., Hong, Z., Liang, R.: 3D printing optics with hybrid material. *Appl. Opt.* **60**, 1809–1813 (2021)
10. Khan, I., Mateus, A., Kamma, C., Mitchell, G.: Part specific applications of additive manufacturing. *Procedia Manuf.* **12**, 89–95 (2017)
11. Ko, K., Jin, S., Lee, S.E., Lee, I., Hong, J.-W.: Bio-inspired bimaterial composites patterned using three-dimensional printing. *Compos. Part B: Eng.* **165**, 594–603 (2019)
12. Möller, T., Trumbore, B.: Fast, minimum storage ray-triangle intersection. *J. Graph. Tools* **2**, 21–28 (1997)
13. Nemhauser, G., Wolsey, L., Fisher, M.: An analysis of approximations for maximizing sub-modular set functions-I. *Math. Program.* **14**, 265–294 (1978)
14. Ngo, T.D., Kashani, A., Imbalzano, G., Nguyen, K., Hui, D.: Additive manufacturing (3D printing): a review of materials, methods, applications and challenges. *Compos. Part B: Eng.* **143**, 172–196 (2018)
15. Peng, X., Kuang, X., Roach, D., Wang, Y., Hamel, C., Lu, C., Qi, H.J.: Integrating digital light processing with direct ink writing for hybrid 3D printing of functional structures and devices. *Addit. Manuf.* **40**, 101911 (2021)
16. Porter, M., Ravikumar, N., Barthelat, F., Martini, R.: 3D-printing and mechanics of bio-inspired articulated and multi-material structures. *J. Mech. Behav. Biomed. Mater.* **73**, 114–126 (2017)
17. Valentinčič, J., Peroša, M., Jerman, M., Sabotin, I., Lebar, A.: Low cost printer for DLP stereolithography. *Strojniški vestnik-J. Mech. Eng.* **63**, 559–566 (2017)
18. Valentine, A., Busbee, T., Boley, J.W., Raney, J., Chortos, A., Kotikian, A., Berrigan, J.D., Durstock, M., Lewis, J.: Hybrid 3D printing of soft electronics. *Adv. Mater.* **29**, 1703817 (2017)

# Supply Chain Resilience: Tactical-Operational Models, a Literature Review



Márcia Batista, João Pires Ribeiro, and Ana Barbosa-Póvoa

**Abstract** The field of Supply Chain Resilience (SCR) has been growing in recent years but has become particularly relevant during the pandemic. Modern supply chains (SC), characterised by their global presence and complex structure, are increasingly vulnerable to disruptions. As a result, companies are shifting from a reactive approach to building resilience. As resilience has been mainly analysed at the strategic level, this work aims to assess the current state of the art on SCR quantitative models that support tactical-operational decisions. This is done through a comprehensive literature review where the SC's main activities were considered, as well as the use of Operations Research (OR) methods to build quantitative approaches. It is also studied how these publications model risk and uncertainty and which resilience metrics have been used since these might be highly influenced by the context in which the models are developed. The results make a set of outcomes available to the community: a set of approaches to deal with the tactical-operational problems; and a list of SCR indicators to measure and monitor SCR. Finally, a research framework for tactical-operational resilient SC is proposed, where the main problems that still require further research attention are identified.

**Keywords** Supply chain resilience · Tactical-operational · Quantitative models · Operations research methods · Metrics

## 1 Introduction

Modern SC have become increasingly exposed to unpredictable events due to globalisation and complex structures involving multiple dispersed nodes (Christopher and Peck [1], Kamalahmadi and Parast [2], Saha et al. [3]). Supply chain risk management struggles to deal with these disruptive events, highlighting the need for resilience management (Fiksel [4]). Recent history has shown that these events will continue

---

M. Batista · J. P. Ribeiro (✉) · A. Barbosa-Póvoa  
CEGIST, Instituto Superior Técnico, Universidade de Lisboa, Lisboa, Portugal  
e-mail: [pires.ribeiro@tecnico.pt](mailto:pires.ribeiro@tecnico.pt)

to occur, making it essential to plan for them to maintain a competitive advantage (Munoz and Dunbar [5]). Although SC faced several disruptive events before, the COVID-19 pandemic (Golan et al. [6], Modgil et al. [7]) and, more recently, the war in Ukraine (Valtonen et al. [8], Faggioni et al. [9]) have further highlighted the need for resilient SC as companies shift their focus from efficiency to resilience (Lund et al. [10]).

Historically, SC have been designed for a narrow range of conditions, leaving them vulnerable to unpredictable events (Alicke and Strigel [11], Christopher and Peck [1]). SC have also become more complex, involving relationships with multiple actors who play vital roles in each other's SC, making it crucial to consider these relationships (Lau et al. [12], Ghosh and Jaillet [13]). The current pandemic has increased awareness of the importance of resilience in SC management (Ponomarov and Holcomb [14], Nguyen et al. [15]). However, existing literature on SCR focuses heavily on qualitative insights, with limited quantitative contributions, especially at the tactical-operational level (Ribeiro and Barbosa-Póvoa [16]). This gap is important to address, as these quantitative models can provide valuable aid to decision-makers in evaluating and adopting strategies towards resiliency.

This work aims to construct a reliable assessment of the state of the art of SCR operations research with a focus on the tactical-operational level through a literature review methodology. Relevant content is retrieved, reflecting the established research questions, which are then thoroughly analysed. Insights that can be drawn from recent research on the effects of the current pandemic are discussed. The focus is then directed to four general categories (distribution, inventory, production, and generic SC). Finally, the metrics used are identified, and the results are discussed, proposing a framework for future research.

The remainder of this work is organised as follows: Sect. 2 describes the methodology adopted throughout the review. Section 3 reviews previous literature in this field of research, positioning and distinguishing this review from existing ones. Following the analysis, a set of research questions are established in Sect. 4 to guide the present work. The material collection procedure is described in Sect. 5, followed by selecting categories to attribute structural dimensions to the sample in Sect. 6. The content of the retrieved papers is then evaluated to answer the research questions in Sect. 7, which are then discussed in Sect. 8 to provide future research paths. Finally, Sect. 9 concludes this paper with final considerations and presents the main outcomes.

## 2 Methodology

The methodology adopted in this work follows an adapted form of the one presented by Ribeiro and Barbosa-Póvoa [16], consisting in six main steps, which will constitute sections of this paper in the following order: Previous literature reviews; Research questions; Material collection; Descriptive analysis; Category selection; Material evaluation.

### 3 Previous Literature Reviews

In order to position the present paper, a search on the Web of Science core collection database was performed for literature reviews on “supply chain” and “resilience” and “review” in the “topic” field, published in English from 2010 (when the number of papers published on the theme increased as presented in Ribeiro and Barbosa-Póvoa [16]) up to December 2021. An initial dataset of 236 results was obtained, then refined to 22 by excluding those that did not specifically address SCR or focus on analysing a particular event.

All authors have acknowledged that research on SCR has been steadily increasing in recent years, with the early 2000s being identified as the starting point and 2003 as a turning point due to the increased focus on developing solutions after highly disruptive events such as the 9/11 attacks, however with scarce publications overall (Hohenstein et al. [17], Ribeiro and Barbosa-Póvoa [16]). The work by Katsaliaki et al. [18] sheds light on the diverse nature of SC disruptions and offers a deeper understanding of their root causes and consequences, providing a comprehensive list of SC disruptions and their mitigation strategies. Despite the increasing attention towards SCR in light of the COVID-19 pandemic, there is still a gap in addressing low-demand products as most papers are regarding high-demand products (Chowdhury et al. [19]).

The literature reviewed has acknowledged that there is a lack of a clear definition of SCR in the literature due to the multidisciplinary nature of the term resilience and the ambiguity surrounding SCR elements and their interdependence. This issue has also been identified as hindering the development of a single definition of SCR (Hosseini et al. [20]). As a result, earlier literature reviews have focused on developing and clarifying conceptual terms such as SCR principles, elements, and strategies, intending to propose well-founded definitions of SCR. The earliest reviews, in 2015, (Hohenstein et al. [17], Tukamuhabwa et al. [21]), where both proposed a solid definition for SCR. The former focused on identifying SCR phases (readiness, response, recovery and growth), elements and metrics, while the latter assessed the most cited strategies. Additional to these reviews, Kamalahmadi and Parast [2] focus on SCR principles, proposing a framework incorporating major components of a resilient SC.

Later on, surged works that pursued other qualitative objectives. Ali et al. [22] took the constructs of the existing SCR definitions to develop a concept mapping framework to provide clarity on the subject of SCR for both managerial and research implications, identifying: phases (pre-disruption, during disruption and post-disruption), strategies and capabilities. Ribeiro and Barbosa-Póvoa [16] propose a framework for a holistic approach towards SCR (Focus event, Performance Level, Adaptive Framing, Speed). Posteriorly, Zavala-Alcívar et al. [23] extended on previous works by presenting three building blocks as a framework, whose goal is to integrate key components for analysis, measurement and management of resilience to enhance sustainable SC.

In fact, there is a three-way attention that still has to be given to SCR. As seen above, conceptualization is still a topic of interest, but academics and practitioners should also devote attention to implementation methods and the measurement of SCR (Negri et al. [24]). By doing that, further relationships with other fields of Supply Chain Management (SCM) can be established, thus positively affecting SC performance.

It can be seen that there are context-specific literature reviews, such as those that propose specific definitions for SCR in third-party logistics providers (Gkanatsas and Krikke [25]) and agri-food supply chains (AFSC) (Stone and Rahimi-fard [26]). Despite some papers presenting quantitative approaches, the depth of these approaches is limited in most cases (Hohenstein et al. [17], Kamalahmadi and Parast [2], Kochan and Nowicki [27], Centobelli et al. [28], Zavala-Alcívar et al. [23]). There is a lack of literature on quantitative models and metrics for SCR (Ali et al. [22], Hosseini et al. [20], Kamalahmadi and Parast [2], Ribeiro and Barbosa-Póvoa [16]). Further research on this topic is important as it can aid decision-makers in adopting effective strategies and assessing performance (Gkanatsas and Krikke [25], Ribeiro and Barbosa-Póvoa [16], Han et al. [29]).

Ribeiro and Barbosa-Póvoa [16] conducted a content analysis to study OR methods and metrics used in SCR models. The authors found that the models were scarce and mostly focused on the strategic level. Hosseini et al. [20] later explored analytical approaches to SCR and clarified the significance of three capacities (absorptive, adaptive and restorative), presenting an analysis of the key drivers for each capacity. The findings were categorized based on the capacities. Han et al. [29] aimed to connect SCR capabilities to performance metrics, identifying 11 capabilities and associating them with three dimensions of SCR, but literature was lacking for three capabilities. The three previous reviews provide a foundation for understanding the integration of resilience into mathematical models within supply chain management (SCM). However, there is still a gap in assessing the key constructs required for tactical-operational models. The current paper aims to reduce this gap; therefore, a set of research questions to support such work is presented in the next section.

## 4 Research Questions

Five research questions were established to attain this paper's goal, which will serve as a guide throughout this work. Considering the recent Covid-19 as a worldwide disruption, and in order to understand how it was treated within SC, the following research question is defined:

1. What insight can be withdrawn on SCR from COVID-19 pandemic responses?

Having the context and focusing on the analysis of SCR quantitative approaches to the tactical-operational problems, the following questions aim to be answered:

2. How has tactical and operational resilience been tackled in SC?
3. How have OR methods been used to support tactical and operational decisions?
4. How have uncertainty and risk been modelled in tactical and operational problems?
5. Which resilience quantitative metrics have been used in tactical and operational problems?

## 5 Material Collection

With the goal of retrieving significant articles, a set of keywords was established to conduct a search on the Web of Science collection database. The keywords considered were “supply chain” and “resilience” with the combination of each of the following terms: “tactical”; “operational”; “quantitative”; “optimization”; “simulation”; “heuristics”; “metrics”; “routing”; “scheduling”; “statistics” and “COVID-19”. All searches were restricted to publications in peer-reviewed journals, written in English, and published between 2010 and December 2021. Further refinements were applied, guaranteeing relevance by meeting the following selection criteria: Articles must have SCR as the main focus; purely qualitative approaches were excluded; articles that focus solely on a strategic-level problem were excluded. After removing the duplicates, a set of 385 papers was obtained, which was reduced to 42 articles, following the rules above. An exception was made to publications addressing the current pandemic, thus retaining six papers for that sole purpose.

## 6 Category Selection

Considering the span of the collected material, it is pertinent to define structural dimensions in order to guide the following material evaluation and appropriately answer the established research questions. Hence, the papers are categorized following the logic depicted in Fig. 1, accounting for three main dimensions: COVID-19 outbreak; Decision level; OR approach. COVID-19 is included because it acts as a disruptive event without precedents; Decision level as this describes in an aggregate form the decisions taken in the SC; OR approach is taken as the basis to build quantitative tools.

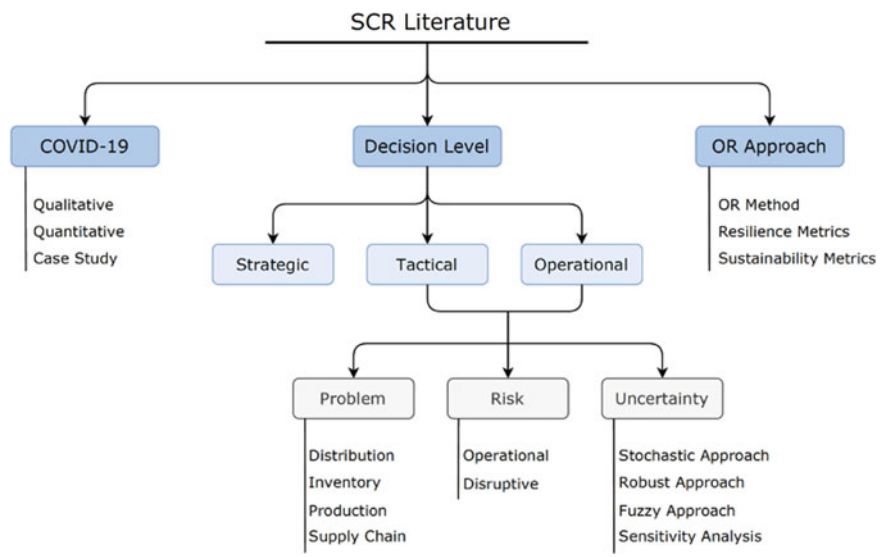


Fig. 1 Mind-map for category selection

7 Material Evaluation

7.1 Insights on SCR from COVID-19 Outbreak Responses (Research Question 1)

On March 11th, 2020, the World Health Organization (WHO) declared COVID-19 a global pandemic, causing worldwide disruptions that lasted longer than expected. Even as the health effects of the pandemic have improved, the crisis’s further evolution and impacts remain uncertain. The severity of this crisis on SC mostly derives from the simultaneous impact on supply and demand due to forces reshaping the business environment, as explored in the previous section. Despite verifying many commonalities across sectors on the caused effects, overall disruptions are essentially industry-specific (Marzantowicz [30]). Hobbs [31], Zhu et al. [32] address food and medical SC, respectively, where a Just-In-Time philosophy is adopted and concluded that these were not able to withstand unpredictable spikes in demand fully, be it from panic buying of groceries or increased necessity for high volumes of personal protective equipment, which led to short-run stockouts.

Marzantowicz [30], Rapaccini [33] report disruptions to SC operations due to protective regulations, decreased orders, and longer transport times. Many managers had to delay deliveries because of difficulties filling and securing transportation. Remko [34] found that companies are reducing reliance on single suppliers and increasing the use of local and nearshore options. de Assunção et al. [35] concluded that shorter SC were less impacted by transport restrictions due to closer proximity to regional

suppliers. Technological investments in digitization and cooperation between entities have been identified as crucial for improving SCR (Remko [34], Marzantowicz [30], Zhu et al. [32]). Inventory management changes and strengthened relationships have also been highlighted as key strategies (Hobbs [31], Zhu et al. [32]). However, more data is needed for further study at the tactical and operational level (Zhu et al. [32]).

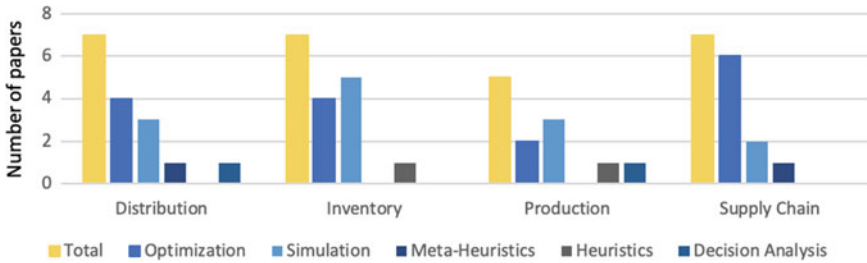
## ***7.2 SCR Quantitative Models on the Tactical and Operational Level (Research Questions 2 and 3)***

In this section, we review studies that address the impact of disruptive events on SC strategies and the use of OR models to facilitate decision-making. Firstly, we examine studies that validate strategies present in qualitative literature and provide insights that can benefit tactical and operational decisions. Secondly, we analyse papers focusing on tactical and operational problems and develop OR models to support decision-making.

Collaboration between entities can be crucial during disruptive events and can aid in increasing the responsiveness of the entire SC. Using agent-based simulation (Lohmer et al. [36]) show that implementing blockchain technology can increase collaboration and agility but require suitable expertise and depend on the length of the disruption. Additionally, Aggarwal and Srivastava [37] use a grey-based Decision Making Trial and Evaluation Laboratory (DEMATEL) approach to study how to build collaborative, resilient SC and find that top management commitment is the most prominent factor. However, they also find that collaboration culture and design resilience in operations strongly influence the remaining success factors.

Various strategies have been proposed to improve SCR in the literature. The work of Azadeh et al. [38] found that redundancy and visibility were effective strategies for a transportation system in a 3-echelon SC. Salehi et al. [39] also found that redundancy was important in a pharmaceutical SC, specifically in the form of additional storage capacity. Zhao et al. [40] studied the impact of reactive and proactive strategies in a complex adaptive system and concluded that proactive strategies were overall superior. Chowdhury and Quaddus [41] identified eight strategies for improving SCR, including backup capacity, building relationships with buyers and suppliers, and adopting technology. Gholami-Zanjani et al. [42]; Nguyen et al. [43]; Chen et al. [44] and Moosavi and Hosseini [45] study the impact of different strategies under SC disruptions, either by MILP or simulation approaches with similar conclusions that profit losses can be reduced by pro-actively looking at SCR. Overall, the literature suggests redundancy, visibility, and proactive strategies can effectively improve SCR.

Models that tackle SCR problems and support decisions at the tactical-operational level are now separated into four general categories: **Distribution problems:** regard coordination of product flows through optimization of supply lead time and rout-



**Fig. 2** OR methods employed in tactical-operational problems

ing assignments; **Inventory problems**: solutions for inventory management decisions; **Production problems**: planning and scheduling considerations for production; **Generic Supply Chain problems**: models that tackle the implementation of recovery policies and timely allocation of resources in more than one SC activity.

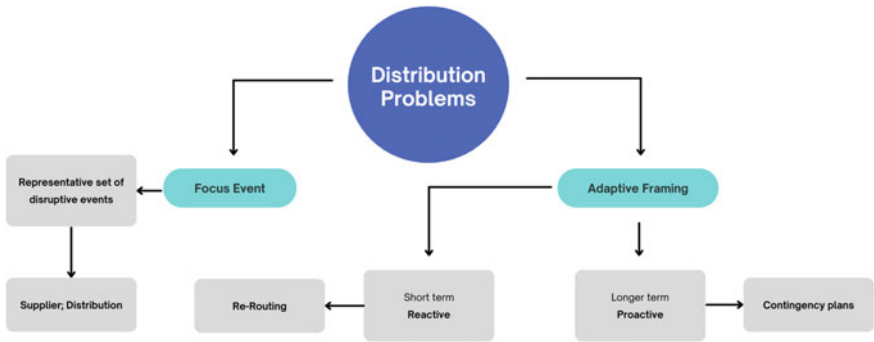
The literature review revealed that production issues are underrepresented in quantitative SCR research, while distribution, inventory, and integrated activities in the recovery process are more prevalent. This is evident in the graph shown in Fig. 2 (yellow bars), which also shows the most commonly used OR method for each problem. Optimization and simulation are the preferred methods, with the occasional use of meta-heuristics, heuristics, and decision-analysis methods.

### 7.2.1 Distribution Problems

In the face of disruptions to a SC, various strategies and methods have been proposed to enhance resilience. Harrison et al. [46] proposed the READI optimization model, which iteratively deletes a component from the network and inserts mitigation strategies, such as re-routing product flow for short-term resilience. Wang et al. [47] focus on disruptions specifically to supplier nodes and investigate how changes in Decision Makers' preferences for suppliers affect the coordination of product flows. Ayoughi et al. [48] propose a four-objective model for a Closed-Loop Supply Chain that balances sustainability and risk minimization, relying on meta-heuristic algorithms for the solution approach.

Studies on SCR to natural hazards can be found for the New York motor-fuel SC and New Zealand's forestry SC. Beheshtian et al. [49] address distribution routing in a bi-stage integer non-linear stochastic model as a planning tool for decision-making in New York. Childerhouse et al. [50] evaluate the New Zealand case through a two-tier modelling approach, using optimization and micro discrete event simulation for port closure scenarios.

A major concern in optimizing product flow is the potential impact of lead time on SC operations. Researchers have used supplier lead time variability as a proxy for SCR, as seen in the work of Colicchia et al. [51], Chang and Lin [52]. These studies test different contingency plans and mitigation strategies to address lead time issues.



**Fig. 3** Proposed generic approach to SCR in distribution problems

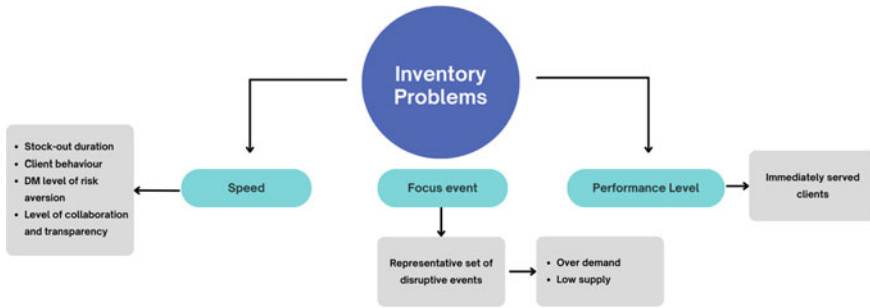
Figure 3 proposes an approach to SCR in distribution problems at the tactical-operational level, highlighting key considerations. Two main approaches to studying SCR are reactive, where re-routing is the most common strategy, and proactive, where representative longer-term contingency plans are implemented.

### Inventory Problems

Unexpected spikes in demand can cause companies to fail to meet customer needs. Schmitt and Singh [53] use the percentage of immediately satisfied customers as a performance metric in their discrete-event simulation to analyse inventory placement and backup capabilities for improved system resilience. Similarly, Wu et al. [54] study the effects of stockouts on consumer response, market share, and stockout duration through agent-based simulation. Lückner and Seifert [55] use stockout quantity and time as SCR metrics to determine optimal risk mitigation inventory levels. To prevent supply-demand mismatches, Spiegler et al. [56] use a simulation model with non-linear control theory to assess the SCR of a grocery SC’s DC replenishment system. Gholami-Zanjani et al. [57] study a Robust Location-Inventory Problem, taking into account the decision-maker’s level of risk aversion.

Recent research on inventory management in SCR has yielded novel approaches, such as the simulation-based optimization model developed by Yang et al. [58], which determines inventory decisions and cost reductions in the event of disruptions to an interconnected network to enhance SCR. Similarly, Gružasuskas et al. [59] examine how information sharing and collaboration can improve forecasting accuracy and market integration. They use agent-based simulation with machine learning algorithms to enhance system resilience.

When dealing with inventory problems in SCR, it is important to consider several factors, as outlined in Fig. 4. The speed linked with SC reaction is crucial to its scope and should include the stock-out duration to examine the trade-off between waiting and investment. The expected client behaviour should also be considered, as



**Fig. 4** Proposed generic approach to SCR in inventory problems

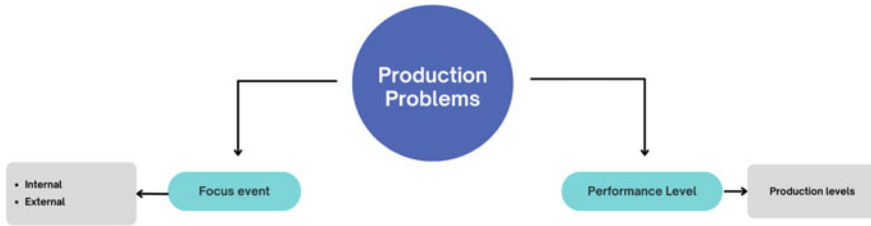
different types of client reactions can have varying consequences. For example, in some industries, not ensuring immediate service may lead to a lost customer, whereas customer loyalty and willingness to wait are higher in others. To study the level of collaboration and transparency in the SC, one should also consider the timespan, as it should allow for complex bullwhip effects to be felt between different entities.

## Production Problems

In addition to inventory management, decision-making systems that can dictate production decisions in response to unexpected events are necessary to enhance SCR. It is important to note that production systems can operate under different strategies, influencing how they must be modelled. For example, Thomas and Mahanty [60] developed a 2-part model for Make-to-Stock environments, where firstly, an analytical investigation based on control engineering techniques and system dynamics. The second part is a simulation of an Inventory and Order-Based Production Control System. For Make-to-Order environments, which involve personalized production, Park et al. [61] propose a novel approach for applying Cyber-Physical Structure and Digital Twins (DT) to SC, using DT simulation to predict and coordinate production of various agents and achieving resilient planning in operation stages in SC production, able to handle disruptions such as the bullwhip effect and ripple effect. DT is also present in Kalaboukas et al. [62] defending that DT are a good option for studying SCR at the tactical/operational level.

In context-specific scenarios, such as the impact of a hurricane on a butanediol SC, Ehlen et al. [63] use a large-scale agent-based simulation to address production scheduling, chemical buying, selling and shipping, and obtain ideal production levels by solving a collection of LP constrained problems.

From Fig. 5, two main elements should be considered to optimize the approach to production problems in SCR: the focus event and the performance level. The focus event can be internal, such as a production line malfunction, or external, such as suppliers not meeting their obligations.



**Fig. 5** Proposed generic approach to SCR in production problems

### Generic Supply Chain Problems

To maintain or acquire a competitive advantage, resilient SCM needs to react quickly and coordinate available resources efficiently during disruptions. Researchers have proposed various models and tools to address this issue. Mao et al. [64] propose a model that optimizes restoration schedules and work crew assignments by a bi-objective non-linear programming model solved by a simulated annealing algorithm. Ivanov et al. [65] develop a dynamic model using optimal program control theory, implemented through simulation, to aid operations and SC planners in adopting reactive recovery policies. Ivanov [66] use attainable sets and optimal control theory to develop two scheduling models for material flow and recovery actions of SC resources. Goldbeck et al. [67] propose a multi-stage stochastic programming model that considers multiple operational adjustments for recovery strategies. Ivanov et al. [68] propose a feedback-driven framework that integrates functional level recovery control with structural recovery control. Ivanov and Dolgui [69] uses discrete-event simulation to investigate efficient recovery policies and redundancy allocation by considering SC overlays of disruptions. Khalili et al. [70] tackle a production-distribution planning problem by developing a two-stage scenario-based mixed stochastic-possibilistic programming model with SCR enhancing options.

### 7.3 Risk and Uncertainty in Tactical and Operational Models (Research Question 4)

When addressing SCR problems, considering risk events, whether operational and/or disruptive, is inevitable. These events are distinguished by their level of impact and frequency, with low-impact and high-frequency events (LIHF) and high-impact and low-frequency ones (HILF) (Khalili et al. [70], Namdar et al. [71]). As expected, disruptive risks are the most considered, given the value of resilience in dealing with HILF events. However, given that these works focus on the tactical-operational level, eight papers were identified that addressed SCR in the face of LIHF events, while five addressed both risks simultaneously.

However, in terms of how these risks are modelled, simple approaches are adopted in the majority of cases. For example, some works that address the occurrence of severe disruptions model such events within a pre-specified time frame (Ehlen et al. [63], Chang and Lin [52], Mao et al. [64]). Other works consider that suppliers are disrupted from the beginning Wang et al. [47] or select which node of the network to be disrupted within a given time length (Harrison et al. [46], Lückert and Seifert [55], Childerhouse et al. [50]).

Another approach is setting different risk profiles that may arise (Schmitt and Singh [53], Wu et al. [54], Beheshtian et al. [49, 72]). These risk profiles are based on interviews with operational personnel and literature (Schmitt and Singh [53]), historical data (Das and Lashkari [73]), or models defined in the extant literature (Yang et al. [58]).

In terms of uncertainty, a few papers have been identified that integrate this condition into their models. Khalili et al. [70] adopt a stochastic approach, taking the production and distribution of new products as imprecise parameters formulated in the form of triangular fuzzy numbers. This approach is also adopted by Ayoughi et al. [48] for dealing with demand, facility costs, and inventory costs parameters, which are taken as uncertain. Goldbeck et al. [67] also adopt a stochastic approach, generating a scenario tree with an input-output method to account for the propagation of risk and interdependency between assets. Lastly, Gholami-Zanjani et al. [57] address uncertainty in selected parameters through a robust approach, using a Monte Carlo method to generate plausible scenarios.

## 7.4 SCR Metrics (Research Question 5)

Relative to resilience metrics that are applied in the context of SCR, it can be said that these are greatly influenced by the end goal of the paper in which they are proposed and/or used, presenting diverse forms. Table 1, where it is represented the main indicators identified in tactical-operational models, confirms this affirmation, where it is possible to verify that there is no unique approach to SCR quantitative models and each application leads to different indicators being used.

Khalili et al. [70] suggest three indicators that are relevant for each activity of the SC, and their weighted sum provides the resilience of the SC as a whole. They compute the Production Resilience as the sum of the initial production capacity and the difference between the decision variables of the additional initial production capacity and the production capacity that is available in a given scenario. Transportation Resilience is calculated based on the difference between the transmission capacity of the transportation mode and the transportation capacity that is available under a scenario. The Inventory Resilience is computed by the difference between the additional capacity of a DC and the emergency inventory level available under a scenario. Nayeri et al. [74] works on inventory location problem with minimum levels of responsiveness as a resilience enabler thus, fostering SCR. Suppliers have

**Table 1** Resilience indicators in tactical-operational models

| Paper                   | Problem addressed | Resilience indicators                        | Formula                                                                                          |
|-------------------------|-------------------|----------------------------------------------|--------------------------------------------------------------------------------------------------|
| Khalili et al. [70]     | Supply chain      | Availability of production capacity (PR)     | $PR_s = \sum_j \sum_t (K_j + w_j - up_{jt}^s)$                                                   |
|                         |                   | Availability of transportation capacity (DR) | $DR_s = \sum_j \sum_c \sum_m \sum_l \sum_t (z_{jc}^{ml} K_m^l - ut_{jct}^s)$                     |
|                         |                   | Availability of emergency inventory (IR)     | $IR_s = \sum_k \sum_c \sum_t (w_{kc} - ub_{kct}^s)$                                              |
| Lücker and Seifert [55] | Inventory         | Stockout surface                             | $S = - \int_{t_1}^{t_2} I(t) dt$                                                                 |
|                         |                   | Mitigation surface                           | $M = \int_0^{t_2} I(t) dt + \int_{t_D}^{\tau} (td)dt + \int_{\tau}^{t_2} t(d + p) dt$            |
| Mao et al. [64]         | Supply chain      | Cumulative loss                              | $R_u = 1 - \frac{\int_{t_e}^{t_e+M} [\varphi(t_0) - \varphi(t)]dt}{M_{max} \times \varphi(t_0)}$ |
|                         |                   | Restoration rapidity                         | $R_m = 1 - \frac{M}{M_{max}}$                                                                    |

to be seen as relevant actors in tactical-operational SCR, as their relationship with the downstream echelons is vital in the short term (Mohammed et al. [75]).

Mao et al. [64] propose normalized indicators in terms of resilience, where the range of its value is between 0 and 1. The resilience of cumulative loss is calculated as the ratio of the area between the SC’s performance and the restoration strategy’s maximum makespan value. Lücker and Seifert [55] construct a resilience metric based on both the stockout quantity and time. The successfully mitigated area and the stockout surface are used as components of the resilience metric. Chang and Lin [52] assess the stability of the net inventory level by observing the time variation between critical inventory points and define impact propagation as the ratio between the stockout duration and the duration of the initial disruptive event.

Li et al. [76] proposed two metrics related to end-customer satisfaction and delivery requirements, while Dixit et al. [77] propose metrics related to order fulfilment and transportation costs. Taghizadeh et al. [78] perform a simulation approach to the impact on SCR in deep-tier SC, with the resilience metric being based on the demand lost. Chen et al. [79] focus on economic interests, proposing three cost-related indicators and one for sales revenue. Gupta et al. [80] uses fuzzy multi-objective linear programming in an Agri-food SC to optimise the operational objective of minimizing costs while maximizing long-term resilience. Zokaei et al. [81] focuses on the reconstruction of SC after a disruption caused by a catastrophe, aiming at providing a balanced solution between different objectives, namely minimizing costs and keeping the high efficiency of the SC. In fact, taking into account costs is an approach many authors follow. However, due to the high impact of disruptions might

lack the trade-off necessary coming from acknowledging the need for profits. Other approaches mitigate that by analysing the Return on Assets and Return on Equity (Azadegan et al. [82], Baghersad and Zobel [83]).

Ahmadian et al. [84] developed a general quantitative model to assess physical networks’ resilience and identify components that require improvements. They measure the resilience of individual components and define criticality as the ability of the network to perform in case of component failure. Sprecher et al. [85] establish resilience indicators for the SC of critical materials, while Sharma et al. [86] focus on the characteristics of trucking companies’ SC. Lastly, Ramezankhani et al. [87] use a hybrid method that combines QFD methodology with DEMATEL to determine the most influential resilience and sustainable factors. Annex 1 lists other papers that present resilience indicators.

8 Discussion and Future Research Directions

The field of SCR is rapidly growing and is expected to continue expanding. We propose a framework for future developments in Fig.6, which encapsulates a future research agenda where three main topics within the development of Tactical-Operational Resilience models should be considered: OR methods; Risk and Uncertainty; Metrics.

Models for tactical-operational decisions in SCR mostly rely on optimization and simulation, while the use of heuristics and decision analysis is under-explored. Future research should focus on both methods, with heuristics particularly valuable for their ability to deliver quick results.

As SCR research advances, common risk management methods are becoming more frequent, but they may not capture the unique aspects of SCR. Further study is needed to understand the impact and frequency of disturbances. Most research has been deterministic, neglecting uncertainty and other parameters. Stochastic approaches should be considered to better reflect real-world dynamics and randomness.

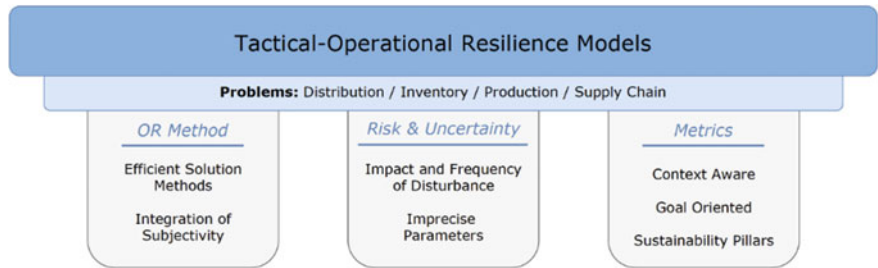


Fig. 6 Research framework on tactical-operational SCR models

Selecting appropriate performance metrics for SCR is critical for accurately reflecting model objectives. Inadequate metrics can lead to inconsistent results, and there is a gap in the literature regarding resilience metrics. Research is needed to bridge this gap by adapting strategies from the strategic level to the tactical-operational level.

## 9 Conclusions

The literature review presented in this paper aims to review the current state of the art on SCR quantitative models, focusing on the tactical and operational decision levels. By conducting a thorough analysis of previous literature reviews, the authors found that SCR is a rapidly growing research field, but quantitative approaches are scarce. To address this gap, the present work focuses on the tactical and operational decision level of SCR.

An overview of the literature developed addressing the COVID-19 pandemic was also executed, and some insights were drawn, although these are limited to early assessments. The developed models were divided into four categories (distribution, inventory, production, SC), and the OR method used within each category was analysed. The category with the least dedicated research was identified as production. The authors then analysed how risk and uncertainty have been modelled and what resilience metrics have been used.

The authors intend to provide useful information that can be used by practitioners and academics alike. To that end, they have identified a list of quantitative resilience indicators, associated mathematical formulas, and the problems where they can be used. They have also provided the most relevant issues to consider when dealing with resilience in specific problems in SCM, which can be used as a checklist for constructing new models that can support decision-makers in building resilience at the tactical and operational levels.

**Acknowledgements** The authors acknowledge the support provided by FCT—Foundation for Science and Technology, I.P., under the project UIDB/00097/2020 and Ph.D. Scholarship: SFRH/BD/148499/2019.

## 10 Annex 1

See Table 2.

**Table 2** Resilience indicators

| Paper                    | Resilience indicators                   |
|--------------------------|-----------------------------------------|
| Munoz and Dunbar [5]     | Recovery                                |
|                          | Performance loss                        |
|                          | Impact                                  |
|                          | Profile Length                          |
| Dixit et al. [77]        | Percentage of unfulfilled demand        |
|                          | Total transportation cost post-disaster |
| Li et al. [76]           | Amount of product delivered             |
|                          | Average delivery distance               |
| Sprecher et al. [85]     | Time lag                                |
|                          | Response speed                          |
|                          | Maximum magnitude                       |
| Sharma and George (2018) | Dimensions of resistive capacity:       |
|                          | Maintenance                             |
|                          | Fuel price variability hedging          |
|                          | Skilled labour and management           |
|                          | Communication and coordination          |
|                          | Security                                |
|                          | Insurance                               |
|                          | Mode flexibility                        |
|                          | Dimensions of restorative capacity:     |
|                          | Risk assessment                         |
|                          | Budget availability                     |
| Ramezankhani et al. [87] | Average inventory                       |
|                          | Economic:                               |
|                          | Cost                                    |
|                          | Part unit profit                        |
|                          | Social:                                 |
|                          | Number of employees                     |
|                          | Employee satisfaction                   |
|                          | Environmental:                          |
|                          | Waste                                   |
|                          | Recyclable waste                        |
|                          | Cost of order loss                      |
| Chen et al. [79]         | Cost of order backlog                   |
|                          | Sales revenue                           |
|                          | Cost of resilience ability              |
|                          |                                         |
| Ahmadian et al. [84]     | Probability of disruption               |
|                          | Impact of the disruption                |
|                          | Recovery to normal state                |
|                          | Criticality                             |
| Li and Zobel (2020)      | Robustness at initial impact            |
|                          | Robustness at full impact               |
|                          | Recovery time                           |
|                          | Average performance retained over time  |

## References

1. Christopher, M., Peck, H.: Building the resilient supply chain. *Int. J. Logist. Manag.* **15**(2), 1–14 (2004)
2. Kamalahmadi, M., Parast, M.M.: A review of the literature on the principles of enterprise and supply chain resilience: major findings and directions for future research. *Int. J. Prod. Econ.* **171**, 116–133 (2016)
3. Saha, A.K., Paul, A., Azeem, A., Paul, S.K.: Mitigating partial-disruption risk: a joint facility location and inventory model considering customers' preferences and the role of substitute products and backorder offers. *Comput. Oper. Res.* (2020)
4. Fiksel, J.: From risk to resilience. In: *Resilient by Design*, pp. 19–34. Springer, Berlin (2015)
5. Munoz, A., Dunbar, M.: On the quantification of operational supply chain resilience. *Int. J. Prod. Res.* **53**(22), 6736–6751 (2015)
6. Golan, M.S., Jernegan, L.H., Linkov, I.: Trends and applications of resilience analytics in supply chain modeling: systematic literature review in the context of the covid-19 pandemic. *Environ. Syst. Decis.* **40**(2), 222–243 (2020)
7. Modgil, S., Gupta, S., Stekelorum, R., Laguir, I.: Ai technologies and their impact on supply chain resilience during covid-19. *Int. J. Phys. Distrib. Logist. Manag.* **52**(2), 130–149 (2021)
8. Valtonen, I., Rautio, S., Lehtonen, J.M.: Designing resilient military logistics with additive manufacturing. *Contin. Resil. Rev.* **5**(1), 1–16 (2023)
9. Faggioni, F., Rossi, M.V., Sestino, A.: Supply chain resilience in the pharmaceutical industry: a qualitative analysis from scholarly and managerial perspectives. *Int. J. Bus. Manag.* **18**, 129 (2023)
10. Lund, S., Manyika, J., Woetzel, J., Barriball, E., Krishnan, M.: Risk, resilience, and rebalancing in global value chains (2020)
11. Aliche, K., Strigel, A.: Supply chain risk management is back, pp. 1–9. McKinsey & Company (2020)
12. Lau, H.C., Agussurja, L., Thangarajoo, R.: Real-time supply chain control via multi-agent adjustable autonomy. *Comput. Oper. Res.* **35**(11), 3452–3464 (2008)
13. Ghosh, S., Jaillet, P.: An iterative security game for computing robust and adaptive network flows. *Comput. Oper. Res.* **138**(105), 558 (2022)
14. Ponomarov, S.Y., Holcomb, M.C.: Understanding the concept of supply chain resilience. *Int. J. Logist. Manag.* **20**(1), 124–143 (2009)
15. Nguyen, T., Li, Z., Spiegler, V., Ieromonachou, P., Lin, Y.: Big data analytics in supply chain management: a state-of-the-art literature review. *Comput. Oper. Res.* **98**, 254–264 (2018)
16. Ribeiro, J.P., Barbosa-Póvoa, A.: Supply chain resilience: definitions and quantitative modelling approaches—a literature review. *Comput. Ind. Eng.* **115**, 109–122 (2018). <https://doi.org/10.1016/j.cie.2017.11.006>
17. Hohenstein, N.O., Feisel, E., Hartmann, E., Giunipero, L.: Research on the phenomenon of supply chain resilience: a systematic review and paths for further investigation. *Int. J. Phys. Distrib. Logist. Manag.* (2015)
18. Katsaliaki, K., Galetsi, P., Kumar, S.: Supply chain disruptions and resilience: a major review and future research agenda. *Ann. Oper. Res.* 1–38 (2021)
19. Chowdhury, P., Paul, S.K., Kaisar, S., Moktadir, M.A.: Covid-19 pandemic related supply chain studies: a systematic review. *Transp. Res. Part E: Logist. Transp. Rev.* **148**(102), 271 (2021)
20. Hosseini, S., Ivanov, D., Dolgui, A.: Review of quantitative methods for supply chain resilience analysis. *Transp. Res. Part E: Logist. Transp. Rev.* **125**, 285–307 (2019). <https://doi.org/10.1016/j.TRE.2019.03.001>
21. Tukamuhbwa, B.R., Stevenson, M., Busby, J., Zorzini, M.: Supply chain resilience: definition, review and theoretical foundations for further study. *Int. J. Prod. Res.* **53**(18), 5592–5623 (2015)

22. Ali, A., Mahfouz, A., Arisha, A.: Analysing supply chain resilience: integrating the constructs in a concept mapping framework via a systematic literature review. *Supply Chain. Manag.: Int. J.* (2017)
23. Zavala-Alcivar, A., Verdecho, M.J., Alfaro-Saiz, J.J.: A conceptual framework to manage resilience and increase sustainability in the supply chain. *Sustainability* (2020)
24. Negri, M., Cagno, E., Colicchia, C., Sarkis, J.: Integrating sustainability and resilience in the supply chain: a systematic literature review and a research agenda. *Bus. Strateg. Environ.* **30**(7), 2858–2886 (2021)
25. Gkanatsas, E., Krikke, H.: Towards a pro-silience framework: a literature review on quantitative modelling of resilient 3pl supply chain network designs. *Sustainability* **12**(10), 4323 (2020)
26. Stone, J., Rahimifard, S.: Resilience in agri-food supply chains: a critical analysis of the literature and synthesis of a novel framework. *Supply Chain. Manag.: Int. J. SCM*–06–2017–0201 (2018). <https://doi.org/10.1108/SCM-06-2017-0201>
27. Kochan, C.G., Nowicki, D.R.: Supply chain resilience: a systematic literature review and typological framework. *Int. J. Phys. Distrib. Logist. Manag.* (2018)
28. Centobelli, P., Cerchione, R., Ertz, M.: Managing supply chain resilience to pursue business and environmental strategies. *Bus. Strateg. Environ.* **29**(3), 1215–1246 (2020)
29. Han, Y., Chong, W.K., Li, D.: A systematic literature review of the capabilities and performance metrics of supply chain resilience. *Int. J. Prod. Res.* 1–26 (2020). <https://doi.org/10.1080/00207543.2020.1785034>
30. Marzantowicz, Ł.: The impact of uncertainty factors on the decision-making process of logistics management. *Processes* **8**(5), 512 (2020)
31. Hobbs, J.E.: Food supply chains during the COVID-19 pandemic. *Can. J. Agric. Econ.* 1–6 (2020). <https://doi.org/10.1111/cjag.12237>
32. Zhu, G., Chou, M.C., Tsai, C.W.: Lessons learned from the covid-19 pandemic exposing the shortcomings of current supply chain operations: a long-term prescriptive offering. *Sustainability* **12**(14), 5858 (2020)
33. Rapaccini, M., Saccani, N., Kowalkowski, C., Paiola, M., Adrodegari, F.: Navigating disruptive crises through service-led growth: the impact of covid-19 on Italian manufacturing firms. *Ind. Mark. Manage.* **88**, 225–237 (2020)
34. Remko, V.H.: Research opportunities for a more resilient post-covid-19 supply chain-closing the gap between research findings and industry practice. *Int. J. Oper. Prod. Manag.* **40**(4), 341–355 (2020)
35. de Assunção, M.V.D., Medeiros, M., Moreira, L.N.R., Paiva, I.V.L., de Souza Paes, D.C.A.: Resilience of the Brazilian supply chains due to the impacts of covid-19. *Holos* **5**, 1–20 (2020)
36. Lohmer, J., Bugert, N., Lasch, R.: Analysis of resilience strategies and ripple effect in blockchain-coordinated supply chains: an agent-based simulation study. *Int. J. Prod. Econ.* **228**(107), 882 (2020)
37. Aggarwal, S., Srivastava, M.K.: A grey-based dematel model for building collaborative resilience in supply chain. *Int. J. Qual. Reliab. Manag.* (2019)
38. Azadeh, A., Atrchin, N., Salehi, V., Shojaei, H.: Modelling and improvement of supply chain with imprecise transportation delays and resilience factors. *Int. J. Log. Res. Appl.* **17**(4), 269–282 (2014)
39. Salehi, V., Salehi, R., Mirzayi, M., Akhavadegan, F.: Performance optimization of pharmaceutical supply chain by a unique resilience engineering and fuzzy mathematical framework. *Hum. Factors Ergon. Manuf. Serv. Ind.* **30**(5), 336–348 (2020)
40. Zhao, K., Zuo, Z., Blackhurst, J.V.: Modelling supply chain adaptation for disruptions: an empirically grounded complex adaptive systems approach. *J. Oper. Manag.* **65**(2), 190–212 (2019)
41. Chowdhury, M.M.H., Quaddus, M.A.: A multiple objective optimization based GFD approach for efficient resilient strategies to mitigate supply chain vulnerabilities: the case of garment industry of Bangladesh. *Omega* **57**, 5–21 (2015)

42. Gholami-Zanjani, S.M., Jabalameli, M.S., Pishvae, M.S.: A resilient-green model for multi-echelon meat supply chain planning. *Comput. Ind. Eng.* **152**(107), 018 (2021)
43. Nguyen, H., Sharkey, T.C., Wheeler, S., Mitchell, J.E., Wallace, W.A.: Towards the development of quantitative resilience indices for multi-echelon assembly supply chains. *Omega* **99**(102), 199 (2021)
44. Chen, J., Wang, H., Zhong, R.Y.: A supply chain disruption recovery strategy considering product change under covid-19. *J. Manuf. Syst.* **60**, 920–927 (2021)
45. Moosavi, J., Hosseini, S.: Simulation-based assessment of supply chain resilience with consideration of recovery strategies in the covid-19 pandemic context. *Comput. Ind. Eng.* **160**(107), 593 (2021)
46. Harrison, T.P., Houm, P., Thomas, D.J., Craighead, C.W.: Supply chain disruptions are inevitable –get readi: resiliency enhancement analysis via deletion and insertion. *Transp. J.* **52**(2), 264–276 (2013)
47. Wang, X., Herty, M., Zhao, L.: Contingent rerouting for enhancing supply chain resilience from supplier behavior perspective. *Int. Trans. Oper. Res.* **23**(4), 775–796 (2016)
48. Ayoughi, H., Dehghani Poteh, H., Raad, A., Talebi, D.: Providing an integrated multi-objective model for closed-loop supply chain under fuzzy conditions with upgral approach. *Int. J. Nonlinear Anal. Appl.* **11**(1), 107–136 (2020)
49. Beheshtian, A., Donaghy, K.P., Geddes, R.R., Rouhani, O.M.: Planning resilient motor-fuel supply chain. *Int. J. Disaster Risk Reduct.* **24**, 312–325 (2017)
50. Childerhouse, P., Al Aqqad, M., Zhou, Q., Bezuidenhout, C.: Network resilience modelling: a New Zealand forestry supply chain case. *Int. J. Logist. Manag.* (2020)
51. Colicchia, C., Dallari, F., Melacini, M.: Increasing supply chain resilience in a global sourcing context. *Prod. Plan. Control* **21**(7), 680–694 (2010)
52. Chang, W.S., Lin, Y.T.: The effect of lead-time on supply chain resilience performance. *Asia Pac. Manag. Rev.* **24**(4), 298–309 (2019)
53. Schmitt, A.J., Singh, M.: A quantitative analysis of disruption risk in a multi-echelon supply chain. *Int. J. Prod. Econ.* **139**(1), 22–32 (2012). <https://doi.org/10.1016/j.ijpe.2012.01.004>
54. Wu, T., Huang, S., Blackhurst, J., Zhang, X., Wang, S.: Supply chain risk management: an agent-based simulation to study the impact of retail stockouts. *IEEE Trans. Eng. Manage.* **60**(4), 676–686 (2013)
55. Lückner, F., Seifert, R.W.: Building up resilience in a pharmaceutical supply chain through inventory, dual sourcing and agility capacity. *Omega* **73**, 114–124 (2017)
56. Spiegler, V., Potter, A.T., Naim, M., Towill, D.R.: The value of nonlinear control theory in investigating the underlying dynamics and resilience of a grocery supply chain. *Int. J. Prod. Res.* **54**(1), 265–286 (2016)
57. Gholami-Zanjani, S.M., Jabalameli, M.S., Klibi, W., Pishvae, M.S.: A robust location-inventory model for food supply chains operating under disruptions with ripple effects. *Int. J. Prod. Res.* **59**(1), 301–324 (2021)
58. Yang, Y., Pan, S., Ballot, E.: Mitigating supply chain disruptions through interconnected logistics services in the physical internet. *Int. J. Prod. Res.* **55**(14), 3970–3983 (2017)
59. Gružas, V., Gimžauskienė, E., Navickas, V.: Forecasting accuracy influence on logistics clusters activities: the case of the food industry. *J. Clean. Prod.* **240**(118), 225 (2019)
60. Thomas, A., Mahanty, B.: Interrelationship among resilience, robustness, and bullwhip effect in an inventory and order based production control system. *Kybernetes* (2019)
61. Park, K.T., Son, Y.H., Noh, S.D.: The architectural framework of a cyber physical logistics system for digital-twin-based supply chain control. *Int. J. Prod. Res.* **59**(19), 5721–5742 (2021)
62. Kalaboukas, K., Rožanec, J., Košmerlj, A., Kiritsis, D., Arampatzis, G.: Implementation of cognitive digital twins in connected and agile supply networks-an operational model. *Appl. Sci.* **11**(9), 4103 (2021)
63. Ehlen, M.A., Sun, A.C., Pepple, M.A., Eidson, E.D., Jones, B.S.: Chemical supply chain modeling for analysis of homeland security events. *Comput. Chem. Eng.* **60**, 102–111 (2014)

64. Mao, X., Lou, X., Yuan, C., Zhou, J.: Resilience-based restoration model for supply chain networks. *Mathematics* **8**(2), 163 (2020)
65. Ivanov, D., Dolgui, A., Sokolov, B., Ivanova, M.: Disruptions in supply chains and recovery policies: state-of-the art review. *IFAC-Papers OnLine* **49**(12), 1436–1441 (2016)
66. Ivanov, D.: Supply chain resilience: Modelling, management, and control. In: *International Series in Operations Research and Management Science*, vol. 265, pp. 45–89. Springer, Cham (2018). [https://doi.org/10.1007/978-3-319-69305-7\\_3](https://doi.org/10.1007/978-3-319-69305-7_3)
67. Goldbeck, N., Angeloudis, P., Ochieng, W.: Optimal supply chain resilience with consideration of failure propagation and repair logistics. *Comput. Chem. Eng.* **133**(101), 830 (2020)
68. Ivanov, D., Dolgui, A., Sokolov, B.: Ripple effect in the supply chain: definitions, frameworks and future research perspectives. In: *International Series in Operations Research and Management Science*, vol. 276, pp. 1–33. Springer, Cham (2019). [https://doi.org/10.1007/978-3-030-14302-2\\_1](https://doi.org/10.1007/978-3-030-14302-2_1)
69. Ivanov, D., Dolgui, A.: Viability of intertwined supply networks: extending the supply chain resilience angles towards survivability: a position paper motivated by COVID-19 outbreak. *Int. J. Prod. Res.* **58**(Forthcoming), 1–12 (2020). <https://doi.org/10.1080/00207543.2020.1750727>
70. Khalili, S.M., Jolai, F., Torabi, S.A.: Integrated production-distribution planning in two-echelon systems: a resilience view. *Int. J. Prod. Res.* **55**(4), 1040–1064 (2017)
71. Namdar, J., Li, X., Sawhney, R., Pradhan, N.: Supply chain resilience for single and multiple sourcing in the presence of disruption risks. *Int. J. Prod. Res.* **56**(6), 2339–2360 (2018)
72. Singh, N.P.: Managing environmental uncertainty for improved firm financial performance: the moderating role of supply chain risk management practices on managerial decision making. *Int. J. Log. Res. Appl.* **23**(3), 270–290 (2020). <https://doi.org/10.1080/13675567.2019.1684462>
73. Das, K., Lashkari, R.: Planning production systems resilience by linking supply chain operational factors. *Oper. Supply Chain. Manag.: Int. J.* **10**(2), 110–129 (2017)
74. Nayeri, S., Tavakoli, M., Tanhaeean, M., Jolai, F.: A robust fuzzy stochastic model for the responsive-resilient inventory-location problem: comparison of metaheuristic algorithms. *Ann. Oper. Res.* 1–41 (2021)
75. Mohammed, A., Naghshineh, B., Spiegler, V., Carvalho, H.: Conceptualising a supply and demand resilience methodology: a hybrid dematel-topsis-possibilistic multi-objective optimization approach. *Comput. Ind. Eng.* **160**(107), 589 (2021)
76. Li, G., Li, L., Zhou, Y., Guan, X.: Capacity restoration in a decentralized assembly system with supply disruption risks. *Int. Trans. Oper. Res.* **24**(4), 763–782 (2017). <https://doi.org/10.1111/itor.12324>
77. Dixit, V., Seshadrinath, N., Tiwari, M.: Performance measures based optimization of supply chain network resilience: a NSGA-ii+ co-kriging approach. *Comput. Ind. Eng.* **93**, 205–214 (2016)
78. Taghizadeh, E., Venkatachalam, S., Chinnam, R.B.: Impact of deep-tier visibility on effective resilience assessment of supply networks. *Int. J. Prod. Econ.* **241**(108), 254 (2021)
79. Chen, L., Dui, H., Zhang, C.: A resilience measure for supply chain systems considering the interruption with the cyber-physical systems. *Reliab. Eng. Syst. Saf.* **199**(106), 869 (2020)
80. Gupta, M., Kaur, H., Singh, S.P.: Multi-echelon agri-food supply chain network design integrating operational and strategic objectives: a case of public distribution system in India. *Ann. Oper. Res.* 1–58 (2021)
81. Zokaee, M., Tavakkoli-Moghaddam, R., Rahimi, Y.: Post-disaster reconstruction supply chain: empirical optimization study. *Autom. Constr.* **129**(103), 811 (2021)
82. Azadegan, A., Modi, S., Lucianetti, L.: Surprising supply chain disruptions: mitigation effects of operational slack and supply redundancy. *Int. J. Prod. Econ.* **240**(108), 218 (2021)
83. Baghersad, M., Zobel, C.W.: Organizational resilience to disruption risks: developing metrics and testing effectiveness of operational strategies. *Risk Anal.* **42**(3), 561–579 (2022)
84. Ahmadian, N., Lim, G.J., Cho, J., Bora, S.: A quantitative approach for assessment and improvement of network resilience. *Reliab. Eng. Syst. Saf.* **200**(106), 977 (2020)

85. Sprecher, B., Daigo, I., Spekkink, W., Vos, M., Kleijn, R., Murakami, S., Kramer, G.J.: Novel indicators for the quantification of resilience in critical material supply chains, with a 2010 rare earth crisis case study. *Environ. Sci. Technol.* **51**(7), 3860–3870 (2017)
86. Sharma, N., Sahay, B.S., Shankar, R., Sarma, P.R.: Supply chain agility: review, classification and synthesis. *Int. J. Log. Res. Appl.* **20**(6), 532–559 (2017). <https://doi.org/10.1080/13675567.2017.1335296>
87. Ramezankhani, M.J., Torabi, S.A., Vahidi, F.: Supply chain performance measurement and evaluation: a mixed sustainability and resilience approach. *Comput. Ind. Eng.* **126**, 531–548 (2018)

# Applying Deep Learning Techniques to Forecast Purchases in the Portuguese National Health Service



José Sequeiros, Filipe R. Ramos, Maria Teresa Pereira, Marisa Oliveira, and Lihki Rubio

**Abstract** Forecasting plays a crucial role in enhancing the efficiency and effectiveness of logistics and supply chain management in the healthcare sector, particularly in financial management within healthcare facilities. Modeling and forecasting techniques serve as valuable tools in this domain, with Artificial Neural Networks (ANN), especially Deep Neural Networks (DNN), emerging as promising options, as indicated by the scientific literature. Furthermore, combining ANN with other methodologies has been a subject of frequent discussion. This study builds on previous research that used historical data to predict expenditure on medicines in Portuguese NHS Hospitals. The focus of this study is to evaluate advantages of Deep Learning methodologies. In addition to traditional approaches as Exponential Smoothing (ES), hybrid models are explored, specifically the combination of DNN with the Box and Jenkins methodology (BJ-DNN). A comparative analysis is conducted to assess the predictive quality and computational cost associated with the forecast models. The findings reveal that ES models provide overall suitability and low computational cost. However, considering the Mean Absolute Percentage Error (MAPE), the BJ-DNN model demonstrates robust forecasting accuracy. In conclusion, this study highlights

---

J. Sequeiros (✉) · M. T. Pereira · M. Oliveira  
School of Engineering of Porto (ISEP), Polytechnic of Porto, Portugal  
e-mail: [1180162@isep.ipp.pt](mailto:1180162@isep.ipp.pt)

M. T. Pereira  
e-mail: [mtp@isep.ipp.pt](mailto:mtp@isep.ipp.pt)

M. Oliveira  
e-mail: [mjo@isep.ipp.pt](mailto:mjo@isep.ipp.pt)

F. R. Ramos  
CEAUL-Centro de Estatística e Aplicações, Faculdade de Ciências, Universidade de Lisboa,  
Lisbon, Portugal  
e-mail: [frramos@ciencias.ulisboa.pt](mailto:frramos@ciencias.ulisboa.pt)

M. T. Pereira · M. Oliveira  
INEGI - Institute of Mechanical Engineering and Industrial Management, Porto, Portugal

L. Rubio  
Department of Mathematics and Statistics, Universidad del Norte, Puerto Colombia, Km. 5 vía  
Puerto Colombia, Barranquilla, Colombia  
e-mail: [lihkir@uninorte.edu.co](mailto:lihkir@uninorte.edu.co)

the potential of Deep Learning methodologies for improving expenditure forecasting in healthcare, while maintaining the favorable attributes of traditional Exponential Smoothing models. These findings contribute to the broader understanding of forecasting techniques in the healthcare sector.

**Keywords** BJ-DNN model · Deep learning · ES models · Forecasting · National health service · Prediction error · Time series

## 1 Introduction

The ability to make informed and accurate decisions is essential for any organization, and the use of statistical techniques and tools plays a crucial role in supporting this process. Forecasting, in particular, is an indispensable tool for predicting future outcomes, and it can be extremely useful in making proactive plans and making informed decisions about production, business, financing, and investments, among other things. Forecasting can contribute to creating a competitive advantage for any entity [1].

The health sector is no exception, and the financial management of its units is an area where accurate forecasting can have significant impacts. With increasing costs in demand for care provision, worsening financial situations, and complicated and time-consuming processes combined with increasing demand, healthcare providers face significant challenges in maintaining efficient and effective services. In this context, the use of forecasting models can be a valuable ally in improving logistics and supply chain management, which ultimately contributes to improving the quality of care provided to patients.

As a result, it is of utmost importance to look for forecasting models that are both accurate, reliable and computationally efficient. Time series forecasting is a popular way to model and predict data that changes over time. It can be used to solve a wide range of problems in the health sector, such as predicting demand for medicines and medical supplies, need for staff, and number of patients.

In this paper, we aim to explore the effectiveness of various time series forecasting methodologies to predict purchases in the Portuguese National Health Service, with a particular focus on deep learning models. Performance of these models will be examined in comparison to traditional methods, specifically evaluating their ability to forecast purchases during periods of abrupt market disruptions and changes. This work proposes new methodologies to the growing body of research on forecasting in the health sector and learn more about the best ways to help with making decisions and managing resources.

## 2 Modeling and Forecasting in Time Series

The study of time series has gained practical interest not only in research but also in practice, with a wide range of applications, allowing future values prediction with some margin of error from a historical record of values [2]. However, time series modeling can be a challenging task due to the dynamic and complex nature of time-dependent data.

Econometric models for time series analysis are based only on historical data to make predictions. The behavior of a time series is characterized by some components, including trend, cyclicity, seasonality, and randomness, where the presence of structural breaks will make the modeling and prediction process difficult [3]. Therefore, identifying underlying patterns and trends in the data is critical to develop accurate forecasting models.

When a time series exhibits a simple dynamic with clear evidence of trend and/or seasonality, classical methodologies, such as exponential smoothing techniques (ES), are often sufficient to provide relatively accurate forecasts with low computational cost [4]. The ease of use and understanding of these methodologies, referring to their strong applicability in the study of time series [5]. In general, these methodologies provide satisfactory results and enable a good understanding of the results. However, as the complexity of the data increases, more sophisticated modeling approaches are required to capture the dynamic and complex nature of the time series data. On the other hand, given the limitations pointed out in the scientific literature for classical methodologies [4, 6], several new approaches in the field of artificial intelligence have contributed a lot to advances in predictive analysis [7]. Under this paradigm, Artificial Neural Networks (ANN), in particular Deep Neural Networks (DNN), have been pointed out in the scientific literature as a very promising option [2, 8, 9]. In addition to the proposal of new ANN architectures, the combination of ANN with other methodologies (hybrid models) has been frequently discussed in the literature in order to improve predictive quality as well as to make the models more efficient [2, 9].

### 2.1 *Exponential Smoothing Models*

Exponential smoothing is a classic method in time series forecasting that has been widely used in many fields, including finance, economics, and healthcare. The ES methodology is based on the explicit modeling of the error, trend, and seasonality in the data, and it has been shown to be effective in generating reliable forecasts with quick implementation and easy understanding, making it a popular choice in the business world [10]. It was first introduced by Brown [11] and Holt [12], and since then, it has been further developed and refined to address different types of time series data. In simpler models, such as Single Exponential Smoothing, the weight given to each historical observation may decrease as observations tend towards the

past, which corresponds to predictive terms. This approach is particularly suitable when data is stationary and future behavior is expected to be similar to past patterns. However, when the data has non-stationary components, such as trend and seasonality, more complex ES models, such as Holt-Winters, can be used to capture these patterns and generate more accurate forecasts. These elaborated models incorporate additional parameters to capture the trend and seasonality components, as well as the level of smoothing for the error term. Equation (1) refers to the Single Exponential Smoothing model, where given weight for each historical observation may decrease as observations tend towards the past (or vice versa) [13]

$$\hat{y}_{t+1|t} = \alpha y_t + (1 - \alpha) \hat{y}_{t|t-1} \quad \forall t \in T \quad (1)$$

where  $\hat{y}_{t+1|t}$  is the forecast for time  $t + 1$  based on information up to time  $t \in T$ , and  $0 \leq \alpha \leq 1$  is the smoothing parameter.

More advanced models include Double Exponential Smoothing and Triple Exponential Smoothing models, which integrate additional components such as trend and/or seasonality.

## 2.2 Deep Neural Networks

The fundamental unit of an artificial neural network (ANN) is the artificial neuron, and these neurons are organized in layers. One notable architecture is the multi-layer perceptron (MLP), which represents an early step towards deep neural networks (DNN). MLPs can be trained using the Backpropagation algorithm [14] and are characterized by the presence of multiple hidden layers composed of artificial neurons. As the field progressed, more complex architectures like recurrent neural networks (RNNs) were developed. In addition to the learning that occurs during each training epoch, RNNs incorporate an additional learning input. To handle the temporal aspect of data, it employs an enhanced backpropagation algorithm known as backpropagation through time (BPTT) [15]. This algorithm enables the to learn long-term dependencies [16]. Long short-term memory (LSTM), a specific type of RNN, has emerged as a powerful variant that can not only learn long-term dependencies but also selectively retain relevant information [17]. LSTMs achieve this by incorporating mechanisms to minimize the cost function based on the importance of recent or older data. While MLPs, RNNs, and LSTMs all fall under the category of deep neural networks, their differences lie primarily at the low level of the individual neurons, shown in Fig. 1. For a more in-depth explanation about DNN models and mathematical background, see [4].

In addition to the mentioned architectures, recent studies have proposed hybrid methodologies (combining DNN with other methodologies). For example is the study by Ramos [18], where a hybrid approach is proposed between the methodology of Box and Jenkins (whose main objective is to 'stabilize' the series – constant means and constant variances over time – before their inclusion in the models [19]) and deep

learning neural networks (DNN), using more robust cross-validation methodologies- than k-fold group. Specifically, by combining DNN architectures (MLP, RNN, and LSTM) with the Box and Jenkins methodology, the authors propose a hybrid model, Box and Jenkins-DNN (BJ-DNN). In addition to remarkable predictive capacity of the proposed model, authors highlight reduction in computational cost (by more than 20%), compared to architectures with memory, such as RNN and LSTM, which, despite good predictive quality, there is a considerable computational cost, in terms of algorithm implementation execution time.

3 Methodological Considerations

In terms of methodological procedures, in this paper, Python 3.7 was used within the Jupyter Notebook environment, and all notebooks are available as open source [20]. To streamline the development process, the code was divided into four distinct steps, each contained in separate *ipynb* files: (1) *ExploratoryDataAnalysis*; (2) *ExponentialTialSmoothing*; (3) *DeepNeuralNetwork*; and (4) *DNNOurApproach*, where BJ-DNN is implemented. All notebooks are available were developed from scratch, based on scientific literature [21, 22].

Regarding the implementation of ES methodologies, nine ES models were considered. Models are described in Table 1, where each one is labeled by a pair of letters, related to the trend component and the seasonality component.

Accordingly, three major groups of models can be identified:

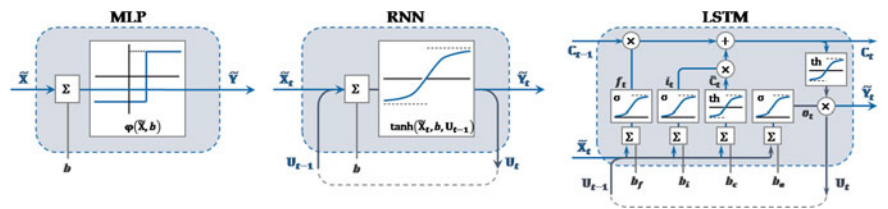


Fig. 1 Comparison of a hidden cells of: MLP, RNN and LSTM (Source Adapted from [4])

Table 1 Classification of ES models

|                                 | Seasonality component (S) |                    |                          |
|---------------------------------|---------------------------|--------------------|--------------------------|
| Trend component (T)             | None (N)                  | Additive model (A) | Multiplicative model (M) |
| None (N)                        | (N, N)                    | (N, A)             | (N, M)                   |
| Additive model (A)              | (A, N)                    | (A, A)             | (A, M)                   |
| Damped additive model ( $A_d$ ) | ( $A_d$ , N)              | ( $A_d$ , A)       | ( $A_d$ , M)             |

Source Adapted from [13]

- Single Exponential Smoothing: (N, N)
- Double Exponential Smoothing: (A, N) e (A<sub>d</sub>, N)
- Triple Exponential Smoothing: (N, A), (A, A), (A<sub>d</sub>, A), (N, M), (A, M) e (A<sub>d</sub>, M)

To select the most adequate ES model, the presence of trend behavior (or trend cycle), seasonality, and the noise present in the observations should be detected. For this, robustness of trend and seasonality were evaluated [23]. Therefore, by representing the trend-cycle component as  $TC_t$ , the seasonal component as  $S_t$ , and the residual (random) component as  $R_t$ , we calculate the measures of trend strength ( $\mathcal{F}_T$ ) and seasonality strength ( $\mathcal{F}_S$ ). Trend strength is defined by (2)

$$\mathcal{F}_T = \max \left( 0, 1 - \frac{\text{Var}(R_t)}{\text{Var}(TC_t + R_t)} \right) \quad (2)$$

where a series with a value of ( $\mathcal{F}_T$ ) close to 1 evidence a strong trend, while in a series without a trend, the value will be 0.

Analogously, the strength of seasonality can be defined by (3)

$$\mathcal{F}_S = \max \left( 0, 1 - \frac{\text{Var}(R_t)}{\text{Var}(S_t + R_t)} \right) \quad (3)$$

where a series with a ( $\mathcal{F}_S$ ) value close to 0 exhibits little or no seasonality, while in a series with strong seasonality, the value will be close to 1.

In addition to determine measures mentioned above, for a more robust choice of the most appropriate ES model for the historical data set, it was decided to consider the information obtained from the Akaike information criterion (AIC) [24] and the Bayesian information criterion (BIC) [25], where the convergence of the parameters of each model for forecasting purposes can be evaluated, thereby choosing the best model (models).

Regarding DNN, in sequential terms, the following steps were taken: (i) Import the data; (ii) Pre-process and/or transform the data; (iii) Define the ANN architecture (MLP, RNN or LSTM) and hyperparameters; (iv) Train and validate the model: the usual cross-validation methodology for time series was used (forward chaining), where the data is split into training samples (60%), validation samples (20%), and test samples (20%); (v) Assess the model: it follows the usual computational procedures in DNN models, through the application of cross-validation methodologies, where the accuracy of the models is evaluated using error metrics during the validation and testing stages (in case it performs poorly, go back to steps (ii) or (iii)); (vi) Forecast.

In addition to the implementation of the routines referring to the BJ-DNN model, only forecast values were obtained from the LSTM model (resulting from the LSTM architecture). According to the literature, LSTM models show better predictive quality in data with disturbances in historical data, compared to models resulting from the MLP and RNN architectures [6, 26]. In each case, a multi-grid was used to explore several possible combinations to define an accurate model.<sup>1</sup>

<sup>1</sup> For more details on implementing DNN models and defining hyperparameters, see [4].

To evaluate a model, it is necessary to compare its forecasted values against the actual price data that the model has not encountered before (test set). This evaluation process generates the forecasting error. The commonly used performance/error metrics are Mean Absolute Error (MAE) and Mean Absolute Percentage Error (MAPE) [27]. Considering the time series  $y_t, t \in T$ , and the past observations from period  $1, \dots, t$ , where  $y_{t+h}$  represents an unknown future value at time  $t+h$  and  $\hat{y}_{t+h}$  denotes its forecast, the prediction error is calculated as the difference between these two values, i.e.,

$$e_{t+h} = y_{t+h} - \hat{y}_{t+h} \quad (4)$$

where  $h$  represents the forecasting horizon.

Additionally, MAE and MAPE are defined, respectively, by

$$MAE = \frac{1}{s} \sum_{i=1}^s |e_{t+i}| \quad (5)$$

$$MAPE = \frac{1}{s} \sum_{i=1}^s \left| \frac{y_{t+i} - \hat{y}_{t+i}}{y_{t+i}} \right| \times 100 \quad (6)$$

where  $s$  corresponds to the number of observations in the forecasting samples (forecasting window).

## 4 Results

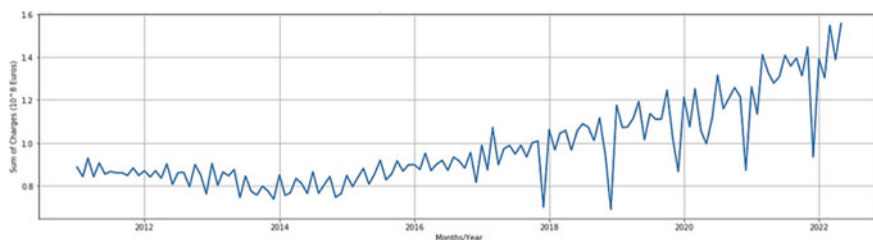
In this research, historical purchasing data from the National Health Service (NHS) was considered, with 137 monthly observations from January 2011 to May 2022.<sup>2</sup> The monthly expenditure on dispensed medicines in public-managed hospital institutions of the National Health Service (NHS) is closely monitored. The data collected pertains to medicines encompassed by the National Hospital Medicines Code (CHNM), which includes both human-use medicines with Marketing Authorization (MA) and medicines authorized for Exceptional Use (EUA) [28].

### 4.1 Time Series Analysis

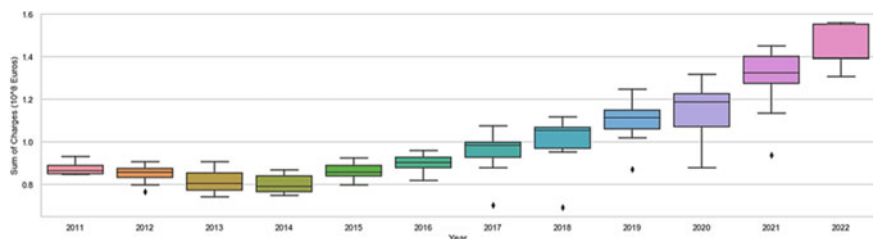
Regarding expenditure on medicines in the Portuguese NHS Hospitals, the historical data is shown in Fig. 2.

---

<sup>2</sup> Data obtained from [28].



**Fig. 2** Expenditure on medicines in the Portuguese NHS Hospitals: graphical representation



**Fig. 3** Expenditure on medicines in the Portuguese NHS Hospitals: graphical representation of annual boxplots

From that chart, it is possible to observe a non-linear pattern with several periods characterized by distinct behaviors. Note the breaks in the last month of the years 2017, 2018, 2019, 2020, and 2021, showing a behavior with some seasonality. This phenomenon can be better discerned when assessing the annual box plots (see Fig. 3), whose referred observations are identified as outliers. It should also be noted that the amplitude of the samples and the interquartile range (IQR) of the box plot that corresponds to the last two years (2020 and 2021) is considerably high, evidencing greater variability in historical data with the occurrence of structural breaks in recent years.

To analyze some features of the time series, Table 2 contains the statistic test and the p-value for the following hypothesis tests: normality tests (Jarque-Bera test and Skewness and Kurtosis tests), stationarity or existence of unit root (ADF test and KPSS test), and independence (BDS test).<sup>3</sup>

As expected, for any significance level, the normality, stationarity, and independence tests are rejected for the time series. There is statistical evidence to: (i) not reject the non-normality of the distribution of the data (with the rejection of the normality test Jarque-Bera); (ii) assume the non-stationarity of the series (due to the null hypothesis not being rejected when doing the ADF test and due to the statistical value corresponding to KPSS being superior to the critical reference values) and (iii) infer about the non-independent and identically distributed (non-iid) since the null

<sup>3</sup> For more details about hypothesis tests, see [4].

hypothesis of the data being independent and identically distributed (iid) has been rejected through the BDS test.

The statistical tests conducted on the time series data indicate that the assumptions of normality, stationarity, and independence are rejected. The evidence suggests that the distribution of the data is non-normal, the series is non-stationary, and the data is non-independent and not identically distributed (non-iid).

## 4.2 Modelling and Forecasting

According to preliminary studies that are consistent with the scientific literature [29], the Charge NHS time series exhibits some statistical properties (non-stationarity, non-linearity, asymmetries, structural breaks) that increase the complexity of modeling and forecasting tasks when using traditional methodologies (e.g., moving average and autoregressive models). The results discussed show some of the limitations in applying the autoregressive models, where the dynamic of the series is not completely captured by these models, which results in large prediction errors. In the present study, in addition to the ES models, DNN models are considered. Concretely, LSTM models and the hybrid BJ-DNN model were used to evaluate their performance in forecasting the time series.

Thus, using data up to December 2021 for modeling, a forecast is made for the first four months of the year 2022 (from January to May). With this, the goal is to figure out how good the predictions are in the short and medium term.

For the ES models, combining the information regarding trend strength ( $\mathcal{F}_T = 0.8536$ ) and seasonality strength ( $\mathcal{F}_T \approx 0.7812$ ) with the information resulting from the AIC and BIC criteria (among the models that exhibited parameter convergence, the candidates for the best models were identified based on the lowest values of the information criteria), two models were selected: (A, M) and ( $A_d$ , M). In Fig 4, you can see the respective graphic representations. In addition to comparing predicted to observed values (forecasting), it was thought useful to represent (with a time window of the last 15 observations) the values adjusted by the model (fitting), from which the accuracy of each model to real data (in-sample forecasts) can be assessed.

**Table 2** Expenditure on medicines in the Portuguese NHS hospitals: normality, stationarity, and independence tests

|           | Normality tests |          |             | Unit root/stationary tests |        | Independence test    |
|-----------|-----------------|----------|-------------|----------------------------|--------|----------------------|
|           | Kurtosis        | Skewness | Jarque-Bera | AADF                       | KPSS   | BDS<br>(Dim.2-Dim.6) |
| Statistic | 1.0149          | 4.4228   | 25.1785     | 5.1107                     | 0.9419 | 14.5761–25.0442      |
| p-value   | 0.3102          | 0.0000*  | 0.0000*     | 1.0000                     | –      | 0.0000*              |

\* $H_0$  is rejected for significance levels of 1%, 5% and 10%

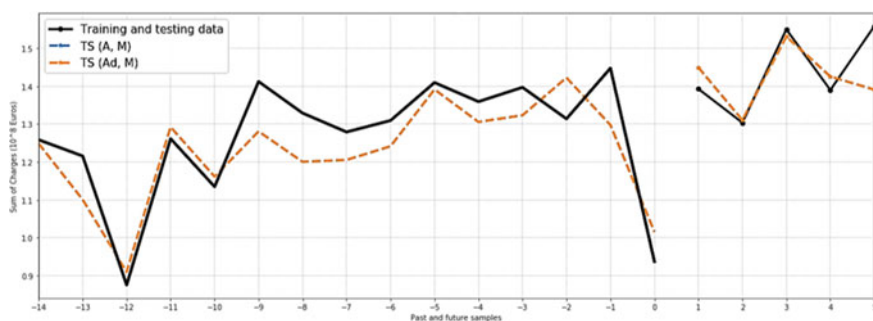
From the graphical analysis, it appears that the two selected models, considering the trend and seasonality components, present similar forecasts (the dashed forecast lines are relatively overlapping). Any of the models (black line) fits the global data well. However, in the medium term, the model starts to lose some predictive quality.

The graphical representations of the DNN models are shown in Fig. 5, (a) LSTM model, and (b) BJ-DNN model. It is possible to observe, for each case: (i) data used to train (shaded area) the model to forecast within the sample (orange line), allowing to understand if the model is overfitting the training data; (ii) data used to train (highlighted in blue horizontal mark) the model to forecast out-of-sample (blue line), allowing the model's performance to be assessed.

From analyzing both figures, it is possible to infer that for both models, the predicted values follow the test data patterns rather accurately. In particular, it is possible to see how BJ-DNN outperforms the LSTM model by comparing their forecasted values (both in and out-of-sample).

### 4.3 Comparing Results

For the DNN models, utilizing LSTM architectures and the BJ-DNN model, the range of MAPE values (lower and upper bounds of 5% trimmed) was calculated for various time horizons: January, February, March, April, and May (refer to Table 3) to conduct a more comprehensive analysis of the out-of-sample forecast.<sup>4</sup> Regarding the ES models, the two selected models present similar MAPE values. An increase in the forecast time horizon does not increase MAPE values. The highest MAPE value in the first forecast month, January, stands out. The high error results from the drop registered in December. Despite the ES models incorporating a seasonality component, this is not enough to be able to capture the dynamic results for January.

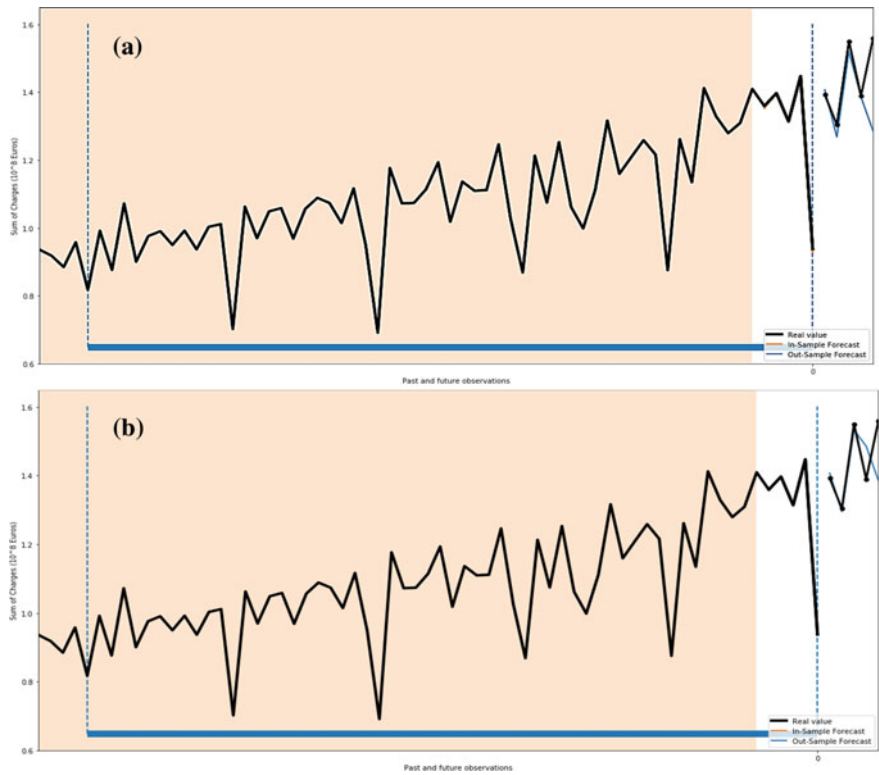


**Fig. 4** ES model (fitting and forecasting)

<sup>4</sup> The parameters of the neural network (weights and bias) benefited from a pseudo-random initialization instead of using a fixed seed [30]. In addition, to avoid outlier results, the forecasting occurred in a loop (60 runs), and the 5% worse and best results were ignored.

In the remaining months, values vary between 1.92% (forecast for March) and 3.82% (forecast for May).

Based on the observed values, the BJ-DNN model has a noticeable general reduction in the amplitude of the MAPE forecast intervals (in each forecast horizon) when compared to the LSTM models in the DNN models. When comparing the three models about forecasting accuracy, despite the instability present in the time series, DNN models seem to be able to better capture the pattern of the series. In the short term (months from January to March), the forecasts resulting from the BJ-DNN model



**Fig. 5** DNN models (fitting and forecasting): **a** LSTM model; **b** BJ-DNN model

**Table 3** Predictions errors of the three DNN models (MAPE)

| Model   | MAPE        |              |           |           |           |
|---------|-------------|--------------|-----------|-----------|-----------|
|         | January (%) | February (%) | March (%) | April (%) | May (%)   |
| ES      | 4.05        | 2.28         | 1.92      | 2.08      | 3.82      |
| LSTM*   | 0.89–1.31   | 1.41–1.93    | 1.72–2.11 | 1.32–1.80 | 3.98–4.43 |
| BJ-DNN* | 0.79–1.12   | 0.68–0.97    | 0.71–1.02 | 1.49–2.01 | 3.68–4.15 |

\*Minimum values—Maximum values (trimmed by 5%) obtained in a total of 60 runs

are better (with an error of around 0.9%). However, for longer forecast horizons (the month of May), the predictive quality of the DNN models decreases. In this case, the forecast values resulting from the classical models are identical or even better.

Therefore, it can be concluded that the models of the chosen methods have different qualities in terms of forecast accuracy. In periods with breaks in historical data, the learning done by the DNN models, particularly the BJ-DNN model, seems to capture the dynamics of the data and produce good forecast values (with low errors). However, at more stable moments in the historical dataset, there seems to be no advantage in implementing more robust methodologies (e.g., neural networks). Classical methodologies, such as ES models, produce identical forecast values (although with greater error) at a lower computational cost.

## 5 Conclusion

Accurate forecasting plays a critical role in efficient healthcare supply chain management and improving clinical outcomes. However, traditional forecasting methods often struggle in complex and rapidly changing environments, particularly when faced with sudden data gaps.

To address these challenges, this study focused on evaluating the effectiveness of deep learning methodologies, specifically the BJ-DNN model, in forecasting purchases within the Portuguese National Health Service. The results demonstrated that the BJ-DNN model, integrating classical procedures, effectively extracted valuable insights from complex and dynamic data. In contrast, traditional methods exhibited limitations in accurately predicting purchases during sudden changes.

While deep learning models like BJ-DNN show promise in healthcare forecasting, it is important to consider their higher computational cost compared to traditional methods. Nevertheless, the cost remains significantly lower than that of the LSTM model. By combining deep learning models with other approaches, healthcare organizations can enhance forecasting accuracy and optimize resource allocation.

Furthermore, applied studies in supply chain management, specifically in forecasting hospital-level medication demand, are crucial to ensure resource efficiency while maintaining high-quality clinical outcomes. The integration of advanced forecasting methodologies, such as the BJ-DNN model, provides valuable tools for healthcare supply chain management. Leveraging the power of deep learning, coupled with traditional forecasting approaches, enables accurate prediction of medication demand, facilitating effective planning, procurement, and distribution strategies.

This research emphasizes the need for integrated forecasting frameworks that harness the strengths of deep learning models, classical approaches, and human expertise. By adopting such frameworks, healthcare organizations can make informed decisions, optimize resource allocation, and enhance inventory management.

Looking ahead, future studies can explore the specific applications of deep learning models in public policy-making and further understand the unique characteristics of healthcare forecasting compared to other time series. These advancements will

contribute to the ongoing improvement of healthcare supply chain management, enabling organizations to achieve resource efficiency without compromising clinical outcomes.

Overall, this study highlights the potential of deep learning methodologies, such as the BJ-DNN model, in healthcare forecasting, supporting informed decision-making and enhancing the efficiency of healthcare supply chains to ultimately benefit patient care.

## References

1. Hang, N.T.: Research on several applicable forecasting techniques in economic analysis, supporting enterprises to decide management. *World Scientific News* **119**, 52–67 (2019)
2. Tealab, A.: Time series forecasting using artificial neural networks methodologies: a systematic review. *Future Comput. Inf. J.* **3**(2) (2020)
3. Pesaran, M.H., Pettenuzzo, D., Timmermann, A.: Forecasting time series subject to multiple structural breaks. *Rev. Econ. Stud.* **73**(4), 1057–1084 (2006)
4. Ramos, F.R.: Data Science na Modelação e Previsão de Séries Económico-financeiras: das Metodologias Clássicas ao Deep Learning. (PhD Thesis, Instituto Universitário de Lisboa - ISCTE Business School, Lisboa, Portugal) (2021)
5. Simões, P., Costa, J., Provenza, M., Xavier, V., Goulart, J.: Modelos de previsão para temperatura e salinidade no fenómeno de ressurgência: Análise dos dados da boia 19 00's34 00'w no período entre 2005 e 2014. In XIX Simpósio de Pesquisa Operacional e Logística da Marinha, Editora Edgard Blucher, Ltda **3**, 1843–1858 (2020)
6. Lopes, D.R., Ramos, F.R., Costa, A., Mendes, D.: Forecasting models for time-series: a comparative study between classical methodologies and deep learning. In: SPE 2021 - XXV Congresso da Sociedade Portuguesa de Estatística. Évora - Portugal (2021)
7. Cavalcante, R.C., Brasileiro, R.C., Souza, V.L.F., Nobrega, J.P., Oliveira, A.L.I.: Computational intelligence and financial markets: a survey and future directions. *Expert Syst. Appl.* **55**, 194–211 (2016)
8. Tkáč, M., Verner, R.: Artificial neural networks in business: two decades of research. *Appl. Soft Comput.* **38**, 788–804 (2016)
9. Sezer, O.B., Gudelek, M.U., Ozbayoglu, A.M.: Financial time series forecasting with deep learning: a systematic literature review: 2005–2019. *Appl. Soft Comput.* **90**, 106–181 (2020)
10. Ord, K.: Charles Holt's report on exponentially weighted moving averages: an introduction and appreciation. *Int. J. Forecast.* **20**(1), 1–3 (2004)
11. Brown, R.G.: *Statistical Forecasting for Inventory Control*. McGraw-Hill, New York (1959)
12. Holt, C.: *Forecasting Seasonals and Trends by Exponentially Weighted Moving Averages*. Carnegie Institute of Technology Graduate School of Industrial Administration, Pittsburgh (1957)
13. Hyndman, R.J., Athanasopoulos, G.: *Forecasting: Principles and Practice* (2nd ed.). OTexts, Melbourne (2018). Retrieved from <https://otexts.com/fpp2/>
14. Rumelhart, D.E., Hinton, G.E., Williams, R.J.: Learning representations by back-propagating errors. *Nature* **323**(6088), 533–536 (1986)
15. Pineda, F.: Generalization of Back propagation to Recurrent and Higher Order Neural Networks. Undefined (1987)
16. Koutník, J., Greff, K., Gomez, F., Schmidhuber, J.: A clockwork RNN. In: *Proceedings of the 31st International Conference on Machine Learning (ICML 2014)*, vol. 5, pp. 3881–3889 (2014)
17. Hochreiter, S., Schmidhuber, J.: Long short-term memory. *Neural Comput.* **9**(8), 1735–1780 (1997)

18. Ramos, F. R., Lopes, D. R., Pratas, T. E.: Deep neural networks: a hybrid approach using box & Jenkins methodology. In: *Innovations in Mechatronics Engineering II*. iCieng 2022. Lecture Notes in Mechanical Engineering, pp. 51–62. Springer, Cham (2022)
19. Box, G., Jenkins, G.M., Reinsel, G.: *Time Series Analysis: Forecasting and Control*, 3rd edn. Prentice Hall, New Jersey (1994)
20. Lopes, D.R., Ramos, F.R.: *Univariate Time Series Forecast* (2020). Retrieved from <https://github.com/DidierRLopes/UnivariateTimeSeriesForecast>
21. Ravichandiran, S.: *Hands-On Deep Learning Algorithms with Python: Master Deep Learning Algorithms with Extensive Math by Implementing Them Using TensorFlow*. Packt Publishing Ltd. (2019)
22. Chollet, F.: *Deep Learning with Python*, 2nd edn. Manning Publications (2021)
23. Wang, X., Smith, K., Hyndman, R.: Characteristic-based clustering for time series data. *Data Min. Knowl. Disc.* **13**(3), 335–364 (2006)
24. Akaike, H.: A new look at the statistical model identification. In: *Selected Papers of Hirotugu Akaike*, pp. 215–222. Springer, New York (1974)
25. Schwarz, G.: Estimating the Dimension of a Model. *Ann. Stat.* **6**(2), 461–464 (1978)
26. Ramos, F.R., Lopes, D.R., Costa, A., Mendes, D.: Explorando o poder da memória das redes neurais LSTM na modelação e previsão do PSI 20. In: *SPE 2021 - XXV Congresso da Sociedade Portuguesa de Estatística*. Évora - Portugal (2021)
27. Willmott, C., Matsuura, K.: Advantages of the mean absolute error (MAE) over the root mean square error (RMSE) in assessing average model performance. *Climate Res.* **30**(1), 79–82 (2005)
28. de Saúde, S.N.: *Despesa com Medicamentos nos Hospitais do SNS* (2023). Retrieved May 24, 2023, from <https://transparencia.sns.gov.pt/explore/dataset/despesa-com-medicamentos-nos-hospitais-do-sns/table/?disjunctive.regiao&sort=tempo>
29. Sequeiros, J.A.: Analysis of the supply chain management in a Portuguese public hospital: A case study [Master's thesis, School of Engineering of Porto (ISEP), Polytechnic of Porto]. Retrieved from <http://hdl.handle.net/10400.22/16915>, (2020)
30. Glorot, X., Bengio, Y.: Understanding the difficulty of training deep feedforward neural networks. In: *Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics (AISTATS 2010)*, Sardinia: JMLR Workshop and Conference Proceedings, pp. 249–256 (2010)

# A Three-Level Decision Support Approach Based on Multi-objective Simulation-Optimization and DEA: A Supply Chain Application



Luís Pedro Gomes , António Vieira , Rui Fragoso , Dora Almeida ,  
Luís Coelho , and José Maia Neves

**Abstract** Decision-making in multi-objective problems is a complex problem with several approaches existing in the literature. Many such approaches typically combine simulation and optimization methods to achieve a robust tool capable of providing useful information to decision-makers. While simulation helps decision-makers to test alternative scenarios and allows uncertainty to be considered, optimization enables them to find the best alternatives for specific conditions. This paper proposes an alternative approach consisting of three levels, wherein we use simulation to model complex scenarios, simulation-optimization to identify the scenarios of the Pareto-front, and Data Envelopment Analysis to identify the most efficient solutions, including those not belonging to the Pareto-front, thereby exploiting the benefits of efficiency analysis and simulation-optimization. This can be useful when decision-makers decide to consider scenarios that, while not optimum, are efficient for their decision-making profile. To test our approach, we applied it to a supply chain design problem. Our results show how our approach can be used to analyze a given system from three different perspectives and that some of the solutions, while not optimum, are efficient. In traditional approaches, such scenarios could be overlooked, despite their efficiency for specific decision-making profiles.

**Keywords** Simulation · Multi-objective simulation-optimization · Data envelopment analysis · Supply chain design · Decision-making · Efficiency

---

L. P. Gomes · A. Vieira (✉) · R. Fragoso · D. Almeida · L. Coelho  
CEFAGE, IFFA—Universidade de Évora, Palácio do Vimioso, Largo Marquês de Marialva, 8,  
7000-809 Évora, Portugal  
e-mail: [avieira@uevora.pt](mailto:avieira@uevora.pt)

D. Almeida  
VALORIZA—Research Center for Endogenous Resource Valorization. Polytechnic Institute of  
Portalegre, Ed. BioBip., Campus Politécnico, 10, 7300-555 Portalegre, Portugal

J. M. Neves  
Centro Algoritmi, Escola de Engenharia—Universidade do Minho, Campus Azurém, 4800-058  
Guimarães, Portugal

# 1 Introduction

Decision-making is an essential process in every business. The decision-makers may require proper methods to help them, especially when dealing with complex systems where the objectives under analysis conflict with themselves, such as in multi-objective complex problems [1]. Some of these methods include the use of optimization and simulation to find solutions that are suitable for decision-makers [2–4].

Simulation is extremely useful and insightful when dealing with the dynamics of systems with considerable variability. It can easily incorporate randomness in any process, enabling the analysis of what-if scenarios. This can lead to valuable insights being obtained, including those triggered by events that were previously not considered or those triggered by propagations throughout the time of certain events, potentially leading to overlooked events considerably affecting the system's overall performance. Even so, simulation is not traditionally used to find optimum solutions. Conversely, optimization can be used to find the best or the best set of solutions, including the set of non-dominated solutions in multi-objective problems [5]. As these methods have pros and cons, many authors combine both in different frameworks. For instance, Wang et al. [3] combined the simulation of the behavior of subway passengers to refine a multi-objective optimization to minimize train operation costs and passenger waiting times. Ramirez et al. [4] used simulation and optimization to minimize oil and gas industry costs and time. Nnene et al. [2] used simulation and optimization to solve a transit network design problem, in which simulation evaluates alternative network solutions by simulating travel demand on them while a multi-objective optimization algorithm searches for efficient network solutions. Simulation-optimization is thus a powerful decision-making tool that has the ability to capture intricate relationships and interactions among several entities in a real-world complex system and identify the best design point [6]. This approach searches design points within the design space to optimize performance measures [7].

The first and most important principle of efficiency is to obtain the best result through the minimum use of resources [8]. Therefore, the success of an organization depends on its efficiency, so the measurement of efficiency helps not only to identify inefficiencies but also helps in the development of the organization through the elimination or minimization of these inefficiencies [8]. Taleb [9], for instance, integrated discrete event simulation and data envelopment analysis (DEA) to measure performance and evaluate the efficiency of potential resource allocation configurations for future performance improvement in emergency departments.

Despite the importance of previous approaches, specific decision-makers may be willing to explore certain dominated or non-optimum scenarios due to individual preferences or because, in their view, there can be other scenarios that, while dominated, are more efficient [10].

Our motivation was to develop a decision-making approach that allows decision-makers to choose the most efficient solution for their specific preferences, in a set of

high uncertainty, even though it is a dominated one.

To test our approach, we selected a hypothetical case study consisting of a supply chain design problem, in which simulation is a well-known effective tool to aid decision-makers [6, 7].

This paper is structured as follows. The next section covers the methodology adopted for this research. Namely, the section briefly covers the proposed approach and describes the problem that was considered to apply the proposed approach. In its turn, the third section presents and discusses the obtained results. Finally, the last section discusses the main conclusions of this research.

## 2 Data and Methodology

This section describes the hypothetical case adopted to apply our approach (the first subsection) and the methodology followed in this research (the second subsection), particularly the three levels of decision-support that were incorporated. This description also details how each of the decision-making levels was modeled.

### 2.1 A Supply Chain Design Problem

This subsection describes the case that was considered for our research. In this sense, we decided to consider the problem of supply chain design, in which several suppliers can be selected to provide materials to a manufacturing company. Supplier selection is of extreme importance, therefore a comprehensive approach to decision-making is highly desirable [11].

The company receives orders modeled with a negative exponential distribution with an average of 15 min, where each order can consume up to nine stock units of a particular material. The company can partner with five suppliers in the market for necessary supplies. The suppliers have proposed to the company the delivery times and fixed costs with orders that are shown in Table 1.

**Table 1** Suppliers data

| Supplier  | Average lead time (days) used in the poisson distribution | Fixed cost per order (€) |
|-----------|-----------------------------------------------------------|--------------------------|
| Supplier1 | 21                                                        | 500                      |
| Supplier2 | 24                                                        | 500                      |
| Supplier3 | 41                                                        | 100                      |
| Supplier4 | 13                                                        | 2000                     |
| Supplier5 | 6                                                         | 5000                     |

Regarding the company costs, it was decided to use a variable rupture cost modeled by the Poisson distribution with an average value of 1500€. We opted for this approach to add increased variability to the model. The company faces the problem of considering several objectives and defining the best supply chain configuration for different decision profiles, namely by considering the following parameters:

- The number of suppliers (there is one parameter for the definition of the existence of each supplier in each simulation scenario);
- Quantity to order;
- Reorder point for the material.

## ***2.2 Proposed Approach: A Three-Level Decision Support Approach***

In a general way, the proposed approach consists, in the first phase, of the modeling and simulation of the case. In the second phase, multi-objective simulation optimization (MOSO) is used to obtain a set of non-dominated solutions based on a decision profile that seeks to maximize some objectives and minimize others. Finally, in the third phase, using Data Envelopment Analysis (DEA), we respond to the decision-maker preferences not incorporated in the second phase. Figure 1 represents a general scheme of the proposed approach.

### **First Level—Modeling and Simulation**

As can be seen, the first level comprises the traditional steps in a simulation approach. Usually, these studies start by defining the problem and the parts that will be modeled and analyzed. In a simulation, when modeling the problem, different approaches can be used, e.g., process modeling or agent modeling. In our case, we used SIMIO (Version 15), which allows the user to use concepts of both approaches in an integrated way [12]. Once the model is developed, it must be validated so the user has confidence in the feedback obtained by the model. In this sense, since we adopted a hypothetical case study, our validation process mainly consisted in varying specific parameters of the model until the desired outcome was reached when analyzing processes, such as variation of stock throughout simulation time, variation of stockouts, and duration of orders to suppliers, among others. If this was an actual case study, this process should be adapted so that the results obtained by the simulation model could match previous observations from the field or data.

As an established method for manufacturing and logistics purposes [13], simulation is particularly well suited to studying complex transportation and logistics systems, with many elements that must be coordinated for the system to function smoothly [14]. Therefore, simulation can significantly promote a multi-decision con-

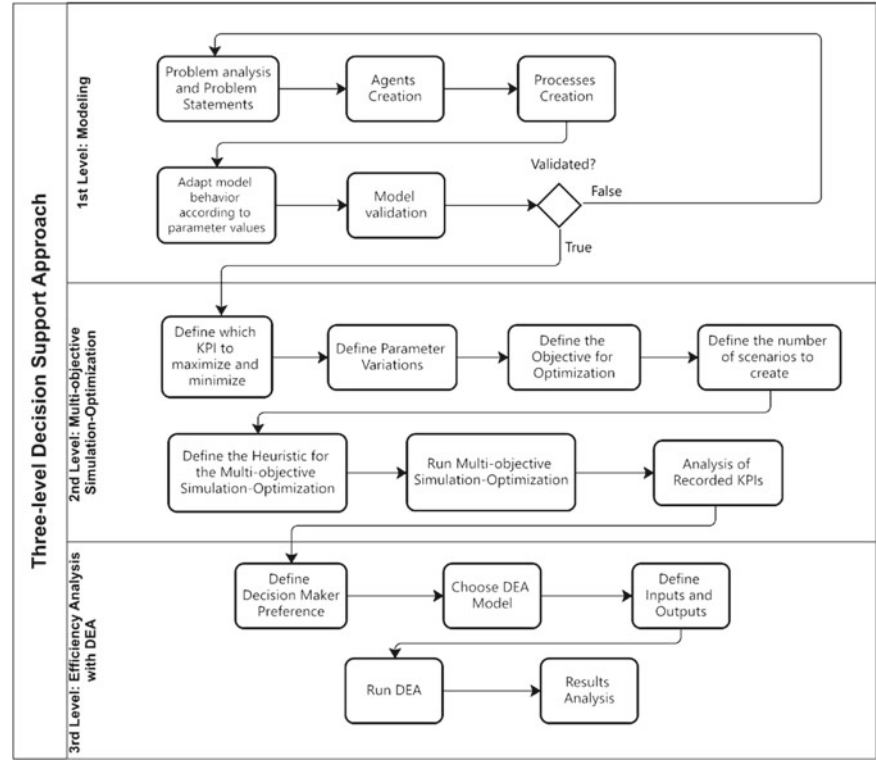


Fig. 1 General scheme of the proposed approach

text that deals with supply chain issues by enabling quantitative assessment of problems originated upstream and downstream of supply chains [15]. The main characteristic of modeling and simulation is its ability to represent business processes of production systems, including, among others, demand forecasts and inventory purchase orders [16, 17]. By simulating a system, managers can test different configurations and identify the most efficient way to operate it [18]. While several processes were defined for the many actions performed in the model by the entities, two main processes can be highlighted. The first considers the internal orders, represented in Fig. 2, and the second consists of the external orders, depicted in Fig. 3.

The internal orders arrive at random intervals and with random material quantities (defined according to the parameters in Table 1. If the material stock is enough, the order is filled, and the quantity supplied is deducted from the existing stock. A material stockout is generated if the amount of stock is insufficient. When the material stock is below the reorder point, an order is generated for the material suppliers.

In the supplier ordering process, a supplier is randomly chosen, and if the selected supplier has an order in transit, another supplier may be chosen, and the order is placed. After this, a random time delay occurs according to the supplier's speci-

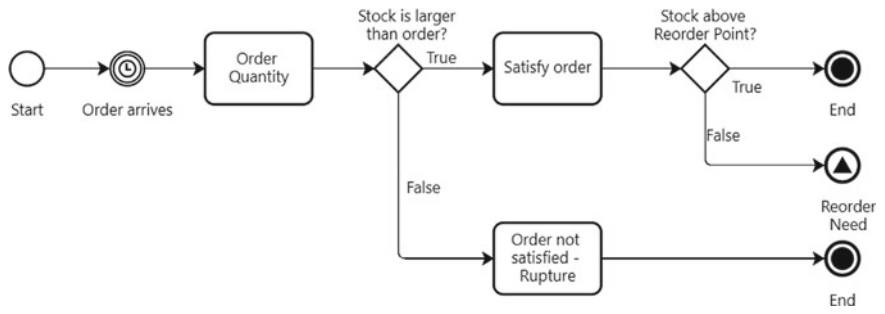


Fig. 2 Process for internal client orders

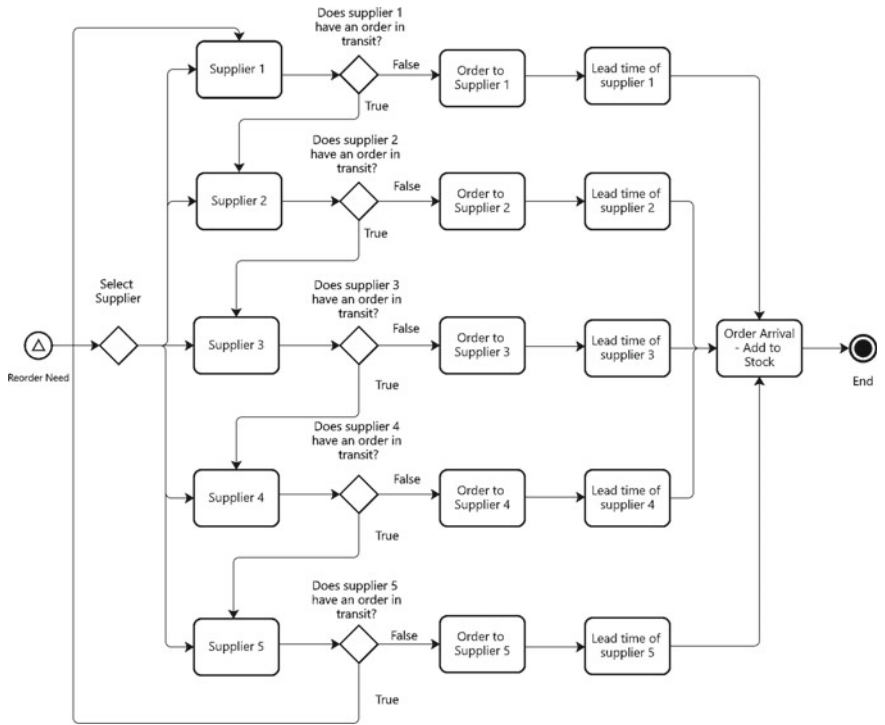


Fig. 3 Process used for orders for suppliers

cations (see Table 1). When the order arrives, the order's quantity is added to the material stock, and the process ends.

## Second Level: Multi-objective Simulation-Optimization

Notwithstanding the traditional benefits of the former level—i.e., using modeling simulation, this is, indeed, an approach with its limitations. This allows users to observe the effects of randomness in selected objectives. However, this does not allow users to determine the best or the best set of scenarios to use under specific established criteria.

Several approaches have been proposed for this purpose, e.g., from the classical stochastic approximation algorithms, response surface methodology, and sample path optimization, to the meta-heuristic approaches including genetic algorithm, simulated annealing, and tabu search (see, for example [19–21] for a survey in this regard). For instance, such approaches may start by limiting the set of possibilities with a method and, after that, applying another method to establish additional performance indicators (when first using optimization and then simulation) or obtain the optimum scenarios after an initial reduction of possibilities (when first using simulation and then optimization). While this is a useful and interesting approach, ours works differently. Instead, the second level does not limit the set of scenarios at the beginning nor the end, i.e., it does not exclude any scenario before simulating a specific number of replications. It evaluates the performance of each scenario and changes the parameters to create scenarios with better performance until the maximum number of scenarios has been created.

In a more detailed way, to use the simulation-optimization approach proposed for the second level, the user has to define the objectives that should be maximized or minimized, as well as the range of values that can be used for the parameters. Thus, each decision-maker profile can establish its preferences. Lastly, the maximum number of scenarios to be created is also defined. Once this setup is done, the simulation optimization engine establishes a set of initial scenarios and runs them for the defined simulation length. Afterward, the obtained objectives are recorded, and a heuristic is used to determine, based on the recorded objectives, the next parameter values to use in the new scenarios that will be created automatically. This process is repeated until the previously defined maximum number of scenarios finishes running. At that point, the scenarios with the best set of parameters and objectives for the considered decision-maker should be obtained, and a Pareto front analysis can be performed with these scenarios, getting a set of non-dominated scenarios.

At this level, the MOSO is used. Real-world optimization problems, such as attaining the lean logistics goal of minimizing inventory while maximizing service level, usually lead to multiple contradictory objectives [5].

The difficulty arises since these give rise to a set of optimal compromise solutions rather than a single optimal solution [5]. As such, it is important to find not just one optimal solution but as many as possible. This way, the decision maker will be better positioned to choose when more compromise solutions are uncovered [5].

The first step consists of defining the objectives that should be maximized and minimized. Supply chain management aims to provide a high service level and simultaneously minimize operating costs [22]. Thus, the main objective considered for our case was the service level, which in this case, is the proportion of orders fulfilled. Regarding the objectives to minimize, the following were considered: total costs and average stock level. Total costs include stock-out costs, holding costs, and ordering costs. The service level is calculated as the percentage of filled orders. Finally, the range of values that can be used for the parameters must also be established. In this case, the parameters included the number of suppliers (from 1 to 5), the rule for their selection (represented in Fig. 3), the quantity to be ordered, and the order point. The maximum number of scenarios to be created was defined as 400. After these steps, the OptQuest ®for SIMIO®is used, in which the simulation-optimization engine establishes a set of initial scenarios and runs them for the defined simulation length. Afterward, the obtained objectives are recorded, and a heuristic is used to determine, based on the recorded objectives, the following parameter values to be used in the new scenarios to be created. This process is repeated until the previously defined maximum number of scenarios has run out. At that point, the obtained scenarios can be submitted to a Pareto front analysis. The service level was defined as the primary response, a minimum of three and a maximum of 50 replications (due to computation availability), a confidence level of 95%, and a relative error of 0.1.

### **Third Level: Efficiency of Dominated and Non-dominated Scenarios with Data Envelopment Analysis**

Obtaining the set of non-dominated scenarios for a given system is an interesting and insightful milestone for decision-makers by itself. However, specific decision-makers may be willing to explore other dominated scenarios since, in their view, there can be different scenarios that, while dominated, are more efficient. Because of this, we decided to complement our approach with the third level, which allows users to evaluate every generated scenario by quantifying their efficiency using DEA. Likewise, in the second level, different decision-maker profiles can be defined using different DEA parameters in the third level.

DEA is aimed at measuring the Decision-Making Units (DMUs) performance according to a predefined decision-maker preference [23]. One of the primary goals of DEA is to measure the efficiency of a DMU through a scalar measure that ranges between zero (the worst performance) and one (the best performance) [24]. One of the pitfalls of DEA is that it does not have a good indicator of model fit [25]. However, this obstacle has been replaced by the satisfaction of a set of theoretical properties: (i) the measure takes values between zero and one; (ii) be monotone; (iii) invariance of the units; (iv) invariance of the translation; and (v) the DMU evaluated is Pareto Koopmans efficient if and only if the measure has a value of one [25].

A DEA analysis should clearly identify what is to be achieved from that analysis [26]. According to Chopra [27], the goal in the design of a supply chain is to structure it in a way that meets customer needs (in this specific case, represented by Service

**Table 2** Inputs and outputs are defined according to the decision maker's different preferences

| Decision maker preference     | Inputs                                         | Outputs                               |
|-------------------------------|------------------------------------------------|---------------------------------------|
| Aversion to high stock levels | Lead time; Total costs;<br>Average stock level | Service level                         |
| Aversion to low service level | Total costs;<br>Average stock level            | Service level                         |
| Aversion to risk exposure     | Total costs;<br>Lead time                      | Service level;<br>Average stock level |

Level, therefore, to be maximized) in a cost-effective manner (being Total Costs, Lead Time, and Average Stock Level representatives of costs, therefore, to be minimized). With that in mind, in this paper, we identify the following three hypothetical possible preferences of a decision-maker:

1. Aversion to high stock levels.
2. Aversion to low service level.
3. Aversion to risk exposure.

The first is related to the efficiency with the minimum stock level. The second concerns the highest level of stock, and finally, the third is associated with a possible aversion to risk (minimizing the possible occurrence of disruptions). Table 2 shows the inputs and outputs defined according to the different profiles of preference of the decision maker.

If both input reduction and output enhancement are desirable goals in a particular application (as in this case, to maximize consumer needs and minimize costs), then a slacks-based measure (SBM) may provide the appropriate model structure to capture a DMU's performance measure [28]. The SBM is an addition to additive models by introducing a measure that makes their efficiency assessment invariant to the units of measure used for the different inputs and outputs [23]. The measure selected was the measure proposed by Tone [29]. This measure satisfies the ensuing properties: (i) it is always between zero and one; (ii) the measure is equal to one if and only if the rated DMU is Pareto-Koopmans efficient; (iii) it is units invariant, and (iv) it is strongly monotonic in inputs and outputs [30]. In the SBM proposed by Tone [29], the K DMU is described by vectors of inputs and outputs, such that:

$$X_{\kappa} = X\lambda + S^{-}; Y_{\kappa} = Y\lambda - S^{+} \quad (1)$$

Where X and Y indicate respectively the matrix of inputs of dimension (m x n) and the matrix of outputs of dimension q1 (h x n), associated with the frontier where the terms S are slacks [31]. Slacks allow non-equality constraints, which indicates that the DMU operates within the frontier, to be expressed by equalities. Efficiency

optimization can thus be considered as an exercise in choosing the slacks and the weights of the components of the vector  $\lambda$ , such that:

$$\min \rho = \frac{1 - \left( \frac{1}{m} \sum_{i=1}^m \frac{S_i^-}{x_{ik}} \right)}{1 - \left( \frac{1}{h} \sum_{r=1}^h \frac{S_r^+}{x_{rk}} \right)} \quad (2)$$

st:

$$\begin{aligned} x_{ik} &= \sum_{j=1}^n x_{ij} \lambda_j + S_i^- \quad i = 1, \dots, m \\ y_{rk} &= \sum_{j=1}^n y_{rj} \lambda_j + S_r^+ \quad r = 1, \dots, m \\ \lambda &\geq 0; \quad S^- \geq 0; \quad S^+ \geq 0 \end{aligned}$$

In this model, the DMU is said to be efficient with a value of unity if, and only if, the DMU is on the frontier of the production possibility set with no input and output slack [32].

### 3 Results

This section presents and discusses the results of applying the proposed approach to the selected problem. In this sense, we show the results that can be obtained at each decision level with our approach in separate subsections.

#### 3.1 Results from the First Level

In terms of decision support, this level allows users to analyze the evolution throughout the simulation time of particular objectives and specific selected performance indicators. This can be helpful for users that already know the parameters they want to experiment with or if they want to understand in more detail what happens in a given scenario. This, indeed, is similar to what most traditional simulation studies can use. Figure 4 shows the run of a simulation experiment.

The first thing that should be highlighted is the simulation optimization's capacity to generate solutions near the Pareto frontier, despite the randomness that occurs throughout each scenario and the wide range of parameters that can be used.

#### 3.2 Results from the Second Level

From the results of the second level, we obtained 360 scenarios with a service level higher than 0,5, of which 21 were non-dominated. Figure 5 shows the results that were obtained.

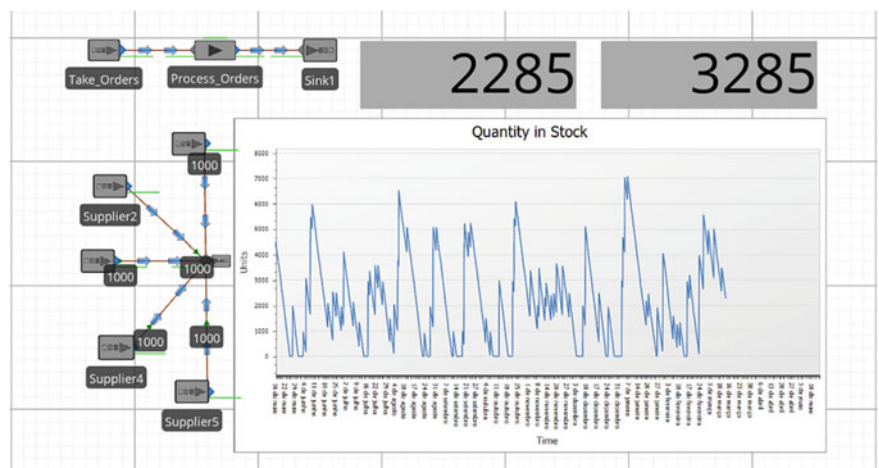


Fig. 4 Run of a simulation experiment

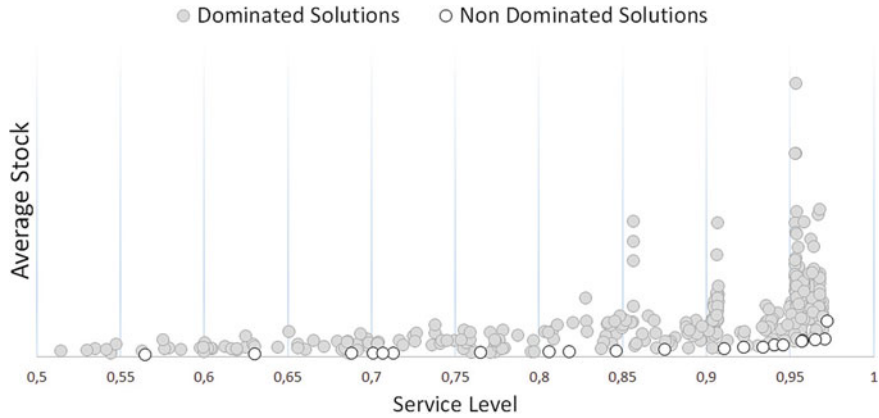
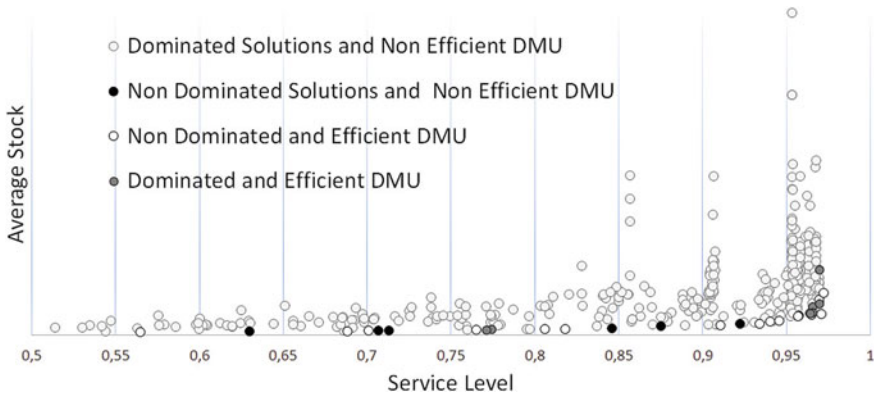


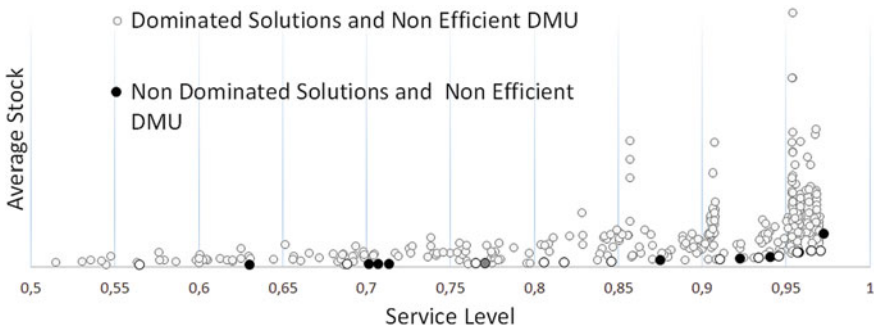
Fig. 5 Sim-Optimization results

3.3 Results from the Third Level

For the decision-maker who has an aversion to stock stocks, the results obtained from the application of DEA allowed to extract of 22 efficient DMUs. Figure 6 graphically represents the comparison between non-dominated scenarios and efficient DMUs, from a decision-maker’s perspective that values stock minimization. From its analysis, it is possible to see that 14 scenarios are simultaneously non-dominated and efficient, and eight dominated scenarios are efficient. These scenarios all exhibit a service level greater than 0.75. Finally, thirteen of the non-dominated scenarios are not efficient from the perspective of a decision-maker that values low stocks. Figure 7 shows the results obtained for the decision-maker with the profile for the aversion



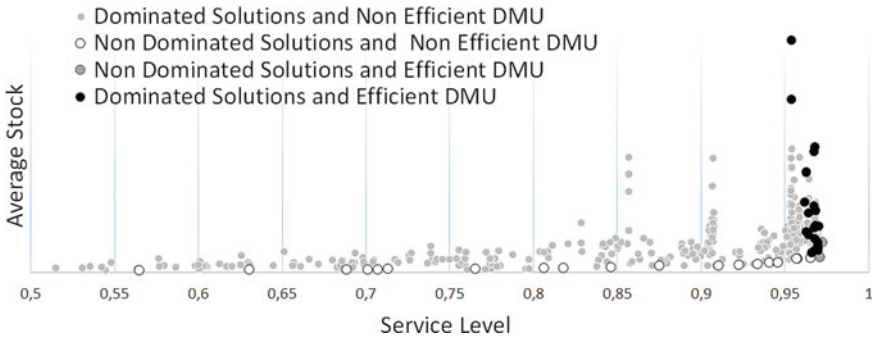
**Fig. 6** Non-dominated, dominated scenarios, and efficient DMUs for stock minimization



**Fig. 7** Non-dominated, dominated scenarios, and efficient DMUs for service level maximization

to low service levels. For the decision-maker profile that values the service level, 12 efficient DMUs were obtained. It is observed from the analysis of Fig. 6 that the vast majority of the efficient DMUs are non-dominated scenarios, with 9 having a service level higher than 0.8. Only one dominated scenario is an efficient DMU. Moreover, a total of eight non-dominated scenarios are not efficient DMUs from the perspective of a decision-maker who values high levels of service. Four have a level of service lower than 0.8. We also highlight that 12 efficient DMUs are also efficient DMUs in the previous decision-making profile that was analyzed, i.e., the aversion to high stock levels. Regarding the third decision-maker profile that was defined, consisting of the aversion to risk exposure, Fig. 8 illustrates the obtained results.

As can be seen, 23 efficient DMUs were obtained. All the efficient DMUs have a service level higher than 0.95. These results can be justified because the average stock level was used as an output for this decision profile. We also highlight the slight overlap between the efficient DMUs in this profile and the efficient solutions, consisting of merely two, as visible in Fig. 7. We also observe little overlap between the efficient DMUs in this profile and the efficient DMUs in the other decision profiles,



**Fig. 8** Non-dominated, dominated scenarios and efficient DMUs for risk minimization

namely six DMUs between the aversion to risk and stock profiles and one between aversion to low service level and risk exposure. In fact, the DMU in question is the only one efficient in the three decision profiles and is simultaneously not dominated.

## 4 Conclusions

This paper proposed an alternative approach for analyzing multi-objective problems, considering uncertainty and efficiency. The approach consists of modeling and simulating a system's dynamics and generating a set of scenarios in pursuit of the Pareto-front. The result is a set of dominated and non-dominated scenarios, which are then analyzed with DEA, considering different decision profiles and both dominated and non-dominated scenarios.

Accordingly, although dominated, we obtained efficient solutions for a given decision-maker profile. Likewise, we also obtained non-dominated solutions that are not efficient for specific decision profiles. We did not discuss the values of the parameters under analysis that originated such scenarios, as that would focus the analysis on the considered problem and not on the approach, which is the scope of this research. Thus, this approach allows the efficiency of solutions to be considered, as well as the non-dominance of others, including the dynamics and uncertainty of complex systems and different decision profiles.

Regardless, our research has some limitations, which should steer our future research endeavors. As a main limitation, we highlight the use of a hypothetical. Real problems may reveal additional interesting challenges and insights into our approach. The selection of suppliers is random, which limits our approach. Different rules could be implemented, including alternative optimization methods or even artificial intelligence, e.g., reinforcement learning. To improve the robustness of the proposed approach, sensitivity analysis should be run in future research.

**Acknowledgements** The authors would like to acknowledge the Portuguese national funds through the FCT—Foundation for Science and Technology, I.P., under the project UIDB/04007/2020.

## References

1. Fragoso, R., Figueira, J.R.: Sustainable supply chain network design: an application to the wine industry in Southern Portugal. *J. Oper. Res. Soc.* **72**, 1236–1251 (2021). <https://doi.org/10.1080/01605682.2020.1718015>
2. Nnene, O.A., Joubert, J.W., Zuidgeest, M.H.P.: A simulation-based optimization approach for designing transit networks. *Public Transp.* (2023). <https://doi.org/10.1007/s12469-022-00312-5>
3. Wang, X., Lv, Y., Sun, H., Xu, G., Qu, Y., Wu, J.: A simulation-based metro train scheduling optimization incorporating multimodal coordination and flexible routing plans. *Transp. Res. Part C Emerg. Technol.* **146**, (2023). <https://doi.org/10.1016/j.trc.2022.103964>
4. Ramírez, A.G., Anguiano, F.I.S.: Simulation based optimization of drilling equipment logistics: a case of study. *Procedia Comput. Sci.* **217**, 866–875 (2023). <https://doi.org/10.1016/j.procs.2022.12.283>
5. Deb, K.: Multi-objective optimization. In: Burke, E.K., Kendall, G. (eds.) *Search Methodologies: Introductory Tutorials in Optimization and Decision Support Techniques*, pp. 403–449. Springer, Boston (2014)
6. Vieira, A., Dias, L.S., Pereira, G.B., Oliveira, J.A., Carvalho, M.S., Martins, P.: Automatic simulation models generation of warehouses with milk runs and pickers. In: *28th European Modeling and Simulation Symposium, EMSS 2016*, pp. 231–241 (2016)
7. Vieira, A.A.C., Dias, L., Santos, M.Y., Pereira, G.A.B., Oliveira, J.: Supply chain risk management: an interactive simulation model in a big data context. *Procedia Manuf.* **42**, 140–145 (2020). <https://doi.org/10.1016/j.promfg.2020.02.035>
8. Panwar, A., Olfati, M., Pant, M., Snasel, V.: A review on the 40 years of existence of data envelopment analysis models: historic development and current trends. *Arch. Comput. Methods Eng.* **29**, 5397–5426 (2022). <https://doi.org/10.1007/s11831-022-09770-3>
9. Taleb, M., Khalid, R., Ramli, R., Nawawi, M.K.M.: An integrated approach of discrete event simulation and a non-radial super efficiency data envelopment analysis for performance evaluation of an emergency department. *Expert. Syst. Appl.* **220**, 119653 (2023). <https://doi.org/10.1016/j.eswa.2023.119653>
10. Kourouxous, T., Bauer, T.: Violations of dominance in decision-making. *Bus. Res.* **12**, 209–239 (2019). <https://doi.org/10.1007/s40685-019-0093-7>
11. Kumar, A., Jain, V., Kumar, S.: A comprehensive environment friendly approach for supplier selection. *Omega (U. K.)* **42**, 109–123 (2014). <https://doi.org/10.1016/j.omega.2013.04.003>
12. Smith, J.S., Sturrock, D.T.: *Simio and Simulation: Modeling, Analysis, Applications*. SIMIO LLC (2021)
13. Straka, M., Lenort, R., Khouri, S., Feliks, J.: Design of large-scale logistics systems using computer simulation hierarchic structure. *Int. J. Simul. Model.* **17**, 105–118 (2018). [https://doi.org/10.2507/IJSIMM17\(1\)422](https://doi.org/10.2507/IJSIMM17(1)422)
14. Ali, R., Khalid, R., Qaiser, S.: A discrete event simulation analysis of the bullwhip effect in a multi-product and multi-echelon supply chain of fast moving consumer goods. *Pak. J. Stat. Oper. Res.* **3**, 561–576 (2020). <https://doi.org/10.18187/pjsor.v16i3.3088>
15. Abideen, A.Z., Binti Mohamad, F.: Empowering supply chain through discrete-event and agent-based simulation—a systematic review and bibliometric analysis. In: *2019 IEEE Conference on Sustainable Utilization and Development in Engineering and Technologies, CSUDET 2019*, pp. 69–74. Institute of Electrical and Electronics Engineers Inc. (2019)

16. Afonso, T., Alves, A., Carneiro, P., Vieira, A.: Simulation pulled by the need to reduce wastes and human effort in an intralogistics project. *Int. J. Ind. Eng. Manag.* **12**, 274–285 (2021). <https://doi.org/10.24867/IJIEEM-2021-4-294>
17. Vieira, A.A.C., Pedro, L., Santos, M.Y., Fernandes, J.M., Dias, L.S.: Data requirements elicitation in big data warehousing. In: Themistocleous, M., Rupino da Cunha, P. (eds.) *Information Systems. EMCIS 2018. Lecture Notes in Business Information Processing*, vol. 341. Springer, Cham (2019). [https://doi.org/10.1007/978-3-030-11395-7\\_10](https://doi.org/10.1007/978-3-030-11395-7_10)
18. Wiśniewski, T., Szymański, R.: Simulation-based optimisation of replenishment policy in supply chains. *Int. J. Logist. Syst. Manag.* **38**, 135 (2021). <https://doi.org/10.1504/IJLSM.2021.113234>
19. Tekin, E., Sabuncuoglu, I.: Simulation optimization: a comprehensive review on theory and applications. *IIE Trans.* **36**, 1067–1081 (2004). <https://doi.org/10.1080/07408170490500654>
20. Fu, M.C.: Feature article: optimization for simulation: theory vs. practice. *INFORMS J Comput.* **14**, 192–215 (2002). <https://doi.org/10.1287/ijoc.14.3.192.113>
21. Fu, M.C., Glover, F.W., April, J.: Simulation optimization: a review, new developments, and applications. In: *Proceedings of the Winter Simulation Conference*, pp. 83–95. IEEE (2005)
22. Peirleitner, A.J., Altendorfer, K., Felberbauer, T.: A simulation approach for multi-stage supply chain optimization to analyze real world transportation effects. In: *2016 Winter Simulation Conference (WSC)*, pp. 2272–2283. IEEE, Washington, DC (2016)
23. Cooper, W.W., Seiford, L.M., Tone, K.: *Data Envelopment Analysis: A Comprehensive Text with Models, Applications, References and DEA-Solver Software*. Springer, New York (2007)
24. Tone, K.: Dealing with desirable inputs in data envelopment analysis: a slacks-based measure approach. *Am. J. Oper. Manag. Inf. Syst.* **6**, 67–74 (2021)
25. Aparicio, J., Monge, J.F.: The generalized range adjusted measure in data envelopment analysis: properties, computational aspects and duality. *Eur. J. Oper. Res.* **302**, 621–632 (2022). <https://doi.org/10.1016/j.ejor.2022.01.001>
26. Cook, W.D., Tone, K., Zhu, J.: Data envelopment analysis: prior to choosing a model. *Omega (Westport)* **44**, 1–4 (2014). <https://doi.org/10.1016/j.omega.2013.09.004>
27. Chopra, S.: *Supply Chain Management: Strategy, Planning, and Operation*. Pearson, Harlow (2019)
28. Cook, W.D., Tone, K., Zhu, J.: Data envelopment analysis: prior to choosing a model. *Omega (U. K.)* **44**, 1–4 (2014). <https://doi.org/10.1016/j.omega.2013.09.004>
29. Tone, K.: A slacks-based measure of efficiency in data envelopment analysis. *Eur. J. Oper. Res.* **130**, 498–509 (2001). [https://doi.org/10.1016/S0377-2217\(99\)00407-5](https://doi.org/10.1016/S0377-2217(99)00407-5)
30. Pastor, J.T., Aparicio, J.: Translation invariance in data envelopment analysis. In: Zhu, J. (ed.) *Data Envelopment Analysis: A Handbook of Models and Methods*, pp. 245–268. Springer, New York (2015)
31. Johnes, G., Tone, K.: The efficiency of higher education institutions in England revisited: comparing alternative measures. *Tert. Educ. Manag.* **23**, 191–205 (2017). <https://doi.org/10.1080/13583883.2016.1203457>
32. Chan, S.G., Karim, M.Z.A., Burton, B., Aktan, B.: Efficiency and risk in commercial banking: empirical evidence from East Asian countries. *Eur. J. Financ.* **20**, 1114–1132 (2014). <https://doi.org/10.1080/1351847X.2012.745008>

# Modelling Forest Fire Spread Through Discrete Event Simulation



Catarina Santos, Ana Raquel Xambre, Andreia Hall, Helena Alvelos, Susete Marques, Isabel Martins, and Filipe Alvelos

**Abstract** Forest fires are becoming a more common occurrence in Portugal as well as worldwide. To extinguish or reduce them more quickly and effectively, it is crucial to understand how they spread. This paper presents a study and a model that shows how wildfires spread, assuming the forest can be represented by a graph, where the nodes correspond to forest stands and the arcs to the path between them. In order to do this, algorithms were developed in Python, using discrete event simulation, that allow modelling the progression of the fire on the graph. This fire propagation model takes into account several aspects of the forest, the wind being the most influential one. Some tests were performed, considering different ignition points, wind directions and wind speeds.

**Keywords** Fire spread · Discrete event simulation · Wind speed · Wind direction

---

C. Santos

Department of Mathematics and Department of Economics, Management, Industrial Engineering and Tourism, University of Aveiro, Aveiro, Portugal  
e-mail: [catarinassantos@ua.pt](mailto:catarinassantos@ua.pt)

A. R. Xambre (✉) · H. Alvelos

Center for Research & Development in Mathematics and Applications and Department of Economics, Management, Industrial Engineering and Tourism, University of Aveiro, Aveiro, Portugal  
e-mail: [raquelx@ua.pt](mailto:raquelx@ua.pt)

H. Alvelos

e-mail: [helena.alvelos@ua.pt](mailto:helena.alvelos@ua.pt)

A. Hall

Center for Research & Development in Mathematics and Applications and Department of Mathematics, University of Aveiro, Aveiro, Portugal  
e-mail: [andreia.hall@ua.pt](mailto:andreia.hall@ua.pt)

S. Marques

Forest Research Centre and School of Agriculture, University of Lisbon, Lisbon, Portugal  
e-mail: [smarques@isa.ulisboa.pt](mailto:smarques@isa.ulisboa.pt)

# 1 Introduction

In recent years, forest fires around the world have been affecting entire communities with the loss of human and animal lives, homes, public infrastructures, air quality, landscapes, among others. The existence of different ecosystems directly interferes with the size, frequency, and intensity of forest fires [1]. Therefore, in order to simulate fire spread, it is necessary to determine the best way to represent the forest and its characteristics. Some of the benefits of using a graph to represent a forest are its considerable small size and the inherent flexibility to model different levels of change, since nodes represent homogeneous sections of terrain and arcs represent the travel path between them [2].

Fuel, weather and topography are three of the factors that most influence how forest fire spreads and the Rothermel's model [3] uses these factors. The model computes the rate,  $R$ , which is the speed of the fire spread, through the following expression [3]:

$$R = \frac{I_R \xi (1 + \phi_w + \phi_s)}{\rho_b \epsilon Q_{ig}} \quad (1)$$

The components in the final equation for the Rothermel surface fire spread model are:

$I_R$ : Intensity of the fire (Btu/ft<sup>2</sup>/min)—Rate of energy release per unit area of fire front

$\xi$ : Propagating flux ratio—Proportion of the reaction intensity that heats adjacent fuel particles to ignition (no wind)

$\phi_w$ : Wind factor—Dimensionless multiplier that considers how wind affects the propagating flux ratio

$\phi_s$ : Slope factor—Dimensionless multiplier that considers how slope affects the propagating flux ratio

$\rho_b$ : Bulk density (lb/ft<sup>3</sup>)—Amount of oven-dry fuel per cubic foot of fuel bed

$\epsilon$ : Effective heating number—Proportion of a fuel particle that reaches ignition temperature at the beginning of flame combustion

$Q_{ig}$ : Heat of preignition (Btu/lb)—Amount of heat required to ignite one pound of fuel.

There are, however, other approaches to model the way fire can propagate across a specific landscape. Some papers simplify the calculation of the rate of fire spread and one of the simple forms is the 10% wind speed rule of thumb for estimating

---

I. Martins

Center of Mathematics, Fundamental Applications and Operations Research and School of Agriculture, University of Lisbon, Lisbon, Portugal  
e-mail: [isabelinha@isa.ulisboa.pt](mailto:isabelinha@isa.ulisboa.pt)

F. Alvelos

ALGORITMI Research Center and LASI, University of Minho, Braga, Portugal  
e-mail: [falvelos@dps.uminho.pt](mailto:falvelos@dps.uminho.pt)

a wildfire's forward rate of spread in forests and shrub lands. The resulting rule of thumb is that the rate of fire spread in forests and shrub lands, under relatively dry conditions, is approximately equal to 10% of the average 10 m open wind speed (5-m average wind speed at a height of 10m above the ground), where both values are expressed in the same units [4]. Other way to model fire spread is considering spread probabilities between cells. The probability of fire spreading from a burning cell to the neighbouring cells was addressed by [5] and can be described as:

$$P_{burn} = P_h(1 + P_{den})(1 + P_{veg})P_w \quad (2)$$

Where  $P_h$  denotes the constant probability that a cell adjacent to a burning cell containing a given type of vegetation and density will catch fire at the next time step;  $P_{den}$ ,  $P_{veg}$  and  $P_w$  are the fire propagation parameters that depend on the density of vegetation, the type of vegetation and the wind speed respectively.  $P_w$  is described by:

$$P_w = e^{C_1 V} e^{V C_2 (\cos \theta - 1)} \quad (3)$$

Where  $V$  denotes the wind speed,  $\theta$  is the angle between the direction of the fire spread and the direction of the wind, and  $C_1$  and  $C_2$  are wind constants. These last two parameters were  $0.045 \text{ m}^{-1} \text{ s}$  and  $0.131 \text{ m}^{-1} \text{ s}$ , respectively. They were determined and defined by [6] by encircling a simulator with a non-linear optimization technique, with the goal of minimizing the discrepancy between the number of burned cells predicted by the simulation and those that were actually burned (respected to a case study). The probability of fire spread increases as the angle between the direction of fire spread and the wind direction decreases because the wind will speed up the spread of the fire in the direction it is blowing in, but will slow it down in directions against the wind [2, 5].

The most common used probability distribution to represent the wind speed is the Weibull distribution since it has been found to fit a wide collection of recorded wind speed data [7]. Therefore, the wind speed probability density function is a two-parameter function (shape and scale) and can be calculated by the following equation [7]:

$$f(v) = \left(\frac{k}{c}\right) \left(\frac{v}{c}\right)^{k-1} \exp \left[ - \left(\frac{v}{c}\right)^k \right], (k > 0, v > 0, c > 1) \quad (4)$$

where  $c$  is the Weibull scale parameter, with units equal to the wind speed units,  $k$  is the unitless Weibull shape parameter and  $v$  is the wind speed [7]. The two parameters of the distribution (shape and scale) can be obtained using known estimation methods such as the maximum likelihood method, the modified maximum likelihood method, and the graphical method. All of them require wind speed data: time-series wind data for the maximum likelihood method, wind speed data in frequency distribution format for the modified maximum likelihood method, and wind speed data in cumulative frequency distribution format for the graphical method [8]. To obtain these distribution parameters in a real situation, it is then necessary to have a sample

of wind speeds in the specific geographic area, in order to be able to apply some estimation method.

The wind intensity can be defined using the Beaufort Scale that classifies the wind intensity in 13 levels [9]. For instance, level 2, described as a light breeze, corresponds to a wind speed within 6 to 11 km/h and level 8, described as a gale, corresponds to a wind speed within 62 to 74 km/h. The first can be associated with light winds and the second with strong winds.

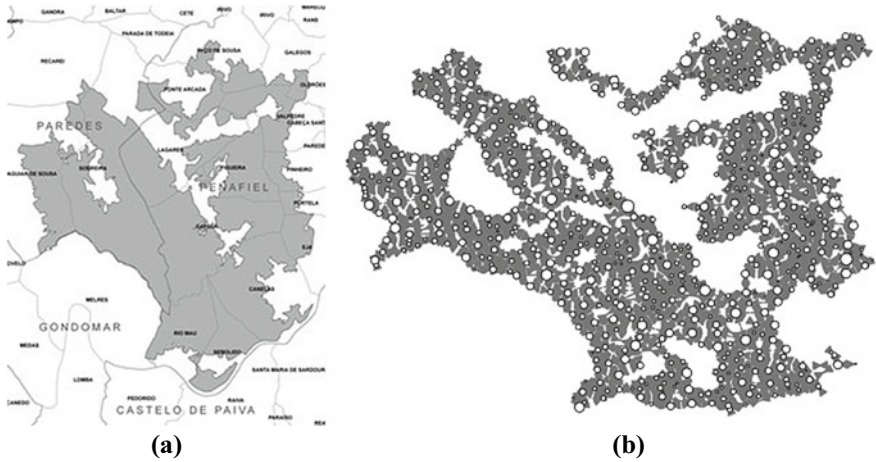
To model a forest fire, one of the first steps is to identify the starting point or points of the fire, i.e. it is necessary to locate the ignition of the fire and analyse whether the fire starts in one or multiple locations. It is possible to determine the probability of ignition in each cell or to model a distribution based on historical ignition locations in specific case studies [10]. Ignition probabilities tend to increase near roads and houses, and with human activity such as smoking, other negligent behaviours, intentional fire setting and other causes, and decrease when associated to land use [11].

Discrete event simulation is a stochastic modelling approach used to model real world systems as a sequence of events in time. It is widely used to address complex and dynamic systems, such as healthcare systems [12], military systems [13] or industrial systems [14]. There are not many studies that use discrete event simulation for modelling the spread of forest fires. Filippi et al. [15] present a novel method that allows the simulation of large-scale/high-resolution systems, focusing on the interface and fire spread. It is based on the discrete event simulation approach, which describes time progression in terms of increments of physical quantities rather than discrete time steps. However, the space representing the forest is not divided into nodes or cells, but into polygons with real coordinates.

The aim of this study is to develop a model to obtain a greater insight into how fire can spread through a specific landscape, and how the factors considered can affect both the behaviour of the fire and the time it takes to burn the entire area. It should thus contribute to the understanding of a problem that highly impacts the environment and the life of the communities.

## 2 Methodology

In this section, the approach applied to the problem at hand is explained. First by detailing how the forest can be represented, followed by the fire ignitions' topic, the fire spread approach and finally the simulation process, where it is possible to understand how the model works with the previously addressed specificities combined.



**Fig. 1** Forest Representation. (a) location of the study area [18]; (b) ZIF Graph derived from (a)

## 2.1 Forest Representation

The forest representation was made through the use of a graph where nodes correspond to stands (contiguous community of trees sufficiently uniform) and arcs to the adjacency between those stands. The nodes' coordinates were obtained by the coordinates of the Forest Intervention Zone (ZIF) of Paiva and Entre Douro e Sousa located in the north of Portugal. The network was then derived from the case study. The existing arcs in the graph were also generated from the information collected, representing the paths that exist between each of the stands. The referred ZIF contains several instances and, although it is possible to apply all these methodologies to any of them, only the northern part of the ZIF will be used as it is the one with the largest area. The graph of the Entre Douro e Sousa ZIF (b), and its real location (a), are represented in Fig. 1, where the white balls represent the stands, and the grey shading represents the arcs. The size of the nodes is proportional to the size of the corresponding stand, so a stand with a significant area becomes perceptible in larger nodes. It is possible to highlight the similarities of the real instance, shown in Fig. 1(a), to the one drawn in Fig. 1(b). This particular graph contains 687 stands which, in turn, covers a total area of approximately 6 611.85 ha.

It is important to have this graph as a representation of the forest since it effectively characterises a real land parcel which, with the specific information obtained from the area, can add value to the mitigation of possible fires and also show the applicability of the present work in real contexts.

**Table 1** Fire ignitions possibilities

| Fire ignitions             | How to obtain the ignition points?             | Simulation                  |
|----------------------------|------------------------------------------------|-----------------------------|
| –Multiple ignitions points | –Chosen                                        | –Variable ignition point(s) |
| –One ignition point        | –Generated by an ignition probability function | –Fixed ignition point(s)    |

**2.2 Fire Ignitions**

Once the network is built it is necessary to determine where and how the fire starts and subsequently spreads. For the fire ignitions the options presented in Table 1 can be assumed.

In wildfires, especially when they are caused by humans, there may be one or more fire ignition points. In this work two ways to obtain these points in the network were assumed: in some situations the ignition points were selected through the analysis of the considered graph, by choosing a node that was in the border of the region, located in a specific area (for example, located in the south area). Alternatively, an ignition probability function was used to randomly choose one or two ignition points. Then, the simulation of the fire spread (that will be addressed in the following section) was implemented and during the run it can be assumed to keep the same ignition point(s) throughout the simulation or change this point(s) each time a new network is run. All this information is presented in the Table 1. Once the fire ignition point(s) are obtained the next step was to simulate how the fire spreads through the network.

**2.3 Fire Spread**

In order to include, in the network, the required information to simulate the fire spread it was decided to assign to each arc of the network two weights: the spread probability between the nodes connected by the arc (*weight1*) and the time of fire spread between those nodes (*weight2*). Assuming that there is fire in the network, the spread probability is the probability of fire spreading from one point to another. The wind direction was used as a way to adjust the probability estimation of fire spread. Each arc of the graph has a certain direction that corresponds to the possible direction of the fire spread and can be seen as a vector, *VecSpread*. Assuming that the wind direction is constant for the whole forest and can be represented as a vector, *VecWind*, it is possible to determine the angle between these two vectors and associate the result obtained in each arc with the respective probability of fire spread. The wind direction was defined considering eight possible scenarios: North, South, East, West, Northeast, Southeast, Northwest and Southwest. This way, the dot product of vectors was applied to each of the arcs of the graph with the wind direction vector. Once the

angle between the wind direction and each of the arcs was obtained, the probability of fire spread was computed according to:

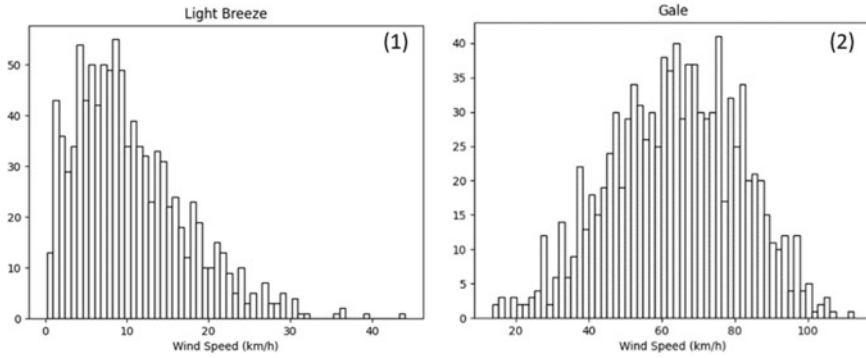
$$p = e^{\cos\theta - 1} * p_h \quad (5)$$

Formula (5) is a simplified version of formula (2), where  $p_h$  denotes the constant probability that a node, adjacent to a burning node, containing a given type of vegetation and density, will catch fire at the next time step. Since there was no available data to determine the value of  $p_h$  for every arc, a constant value of 0.85 was assumed for the entire graph, after some initial tests, in order to obtain realistic values regarding the burnt area. Moreover,  $\theta$  is the angle between the wind and the fire spread directions and it can be noted that the probability of propagation ( $p$ ) is greater when the angle between the arc and the wind direction is small, and is equal to  $p_h$  when the angle between the wind direction and the arc is zero.

The fire spread probability is one of the weights associated with each arc but, in addition to this parameter, there is another weight: the time of fire spread, i.e., the time it takes for the fire to spread from one stand to another. It is possible to obtain this time if the distance travelled and the speed of fire spread are known. The distance between nodes is obtained through the Euclidean distance. The speed of fire spread is a more complicated measurement to obtain accurately. It was decided to use two different ways to calculate this measure: (i) the Rothermel model, and (ii) the 10% wind speed rule of thumb.

In these two models (Rothermel and 10% wind speed) it is always necessary to consider the wind speed as an input. The wind speed was defined by the Weibull distribution with two parameters: scale and shape. Each of the parameters were obtained considering the results of the wind speed that was effectively being considered. Thus, two wind scenarios defined by the Beaufort scale were simulated: level 2 of the Beaufort scale, light breeze with winds between 6 and 11 km/h; and level 8 of the Beaufort scale, gale with winds between 62 and 74 km/h. For light breeze, the shape and scale parameters considered were 1.6 and 1, respectively. For gale, the shape and scale parameter considered were 4 and 6, respectively. These shape and scale parameters used were obtained by testing different values until the observations for the wind speed were compatible with the abovementioned intervals of speed, for light breeze and for gale (see in Fig. 2 the histograms of a thousand wind speed observations). The 95% confidence intervals for the mean of these observations were approximately [10, 11] km/h for light breeze, and [63, 65] km/h for gale. The Rothermel model was used in order to obtain the rate of fire spread that can be seen as the speed of fire spread itself. However, the only component that did not remain constant in the model was the wind speed. Therefore, as the other parameters remained constant, and the only direct influence on the speed of fire spread included was the wind, the values of all the other parameters of the Rothermel model were obtained from the literature [3], and the final equation was as follows:

$$R = \frac{I_R \xi (1 + \phi_w + \phi_s)}{\rho_b \in Q_{ig}} \quad (6)$$



**Fig. 2** Histograms of a thousand wind speed observations given by the Weibull distribution for (1) light breeze and (2) gale

$$I_R = 4.4967e^{03} BTU/ft^2min$$

$$\xi = 1.4782e^{-01}$$

$$\phi_w \text{ is not constant}$$

$$\phi_s = 3.8930e^{-01}$$

$$\rho_b = 1.8119e^{00} lb/ft^3$$

$$\epsilon = 9.2223e^{-01}$$

$$Q_{ig} = 3.9435e^{02} BTU/lb.$$

The  $\phi_w$ , which represents the wind factor, is the only parameter in the expression that varies since it depends on the wind speed, and it is determined according to the Rothermel model [3].

Alternatively, the 10% wind speed rule of thumb gives the speed of fire spread, under relatively dry conditions, and is approximately equal to 10% of the wind speed, where both values are expressed in the same units.

In the next section the necessary steps for simulating fire spread will be explained.

## 2.4 Discrete Event Simulation

The fire propagation along the network was defined based on discrete event simulation (DES). Given a graph with nodes and arcs that represents the forest, each arc has two distinct weights, as explained in the previous section: *weight1* represents the spread probability between the nodes, and *weight2* represents the time of fire spread between the respective nodes.

Before starting the simulation process, it is necessary to choose the ignition node, i.e., the place where the fire starts. It is also required to create a priority queue that will contain the active nodes along the propagation. The active nodes are defined as those that the fire can reach at a certain timestamp. The queue is arranged in ascending order

according to the fire propagation times of the nodes and is processed by removing nodes from it as they change from the 'burning' state to the 'burned' state. Therefore, this priority queue contains the next nodes that the fire can reach and the time it takes to get there. The process starts with a fire ignition that corresponds to the initial node being added to the priority queue with a null propagation time. While the queue is not empty, the simulation process does not finish.

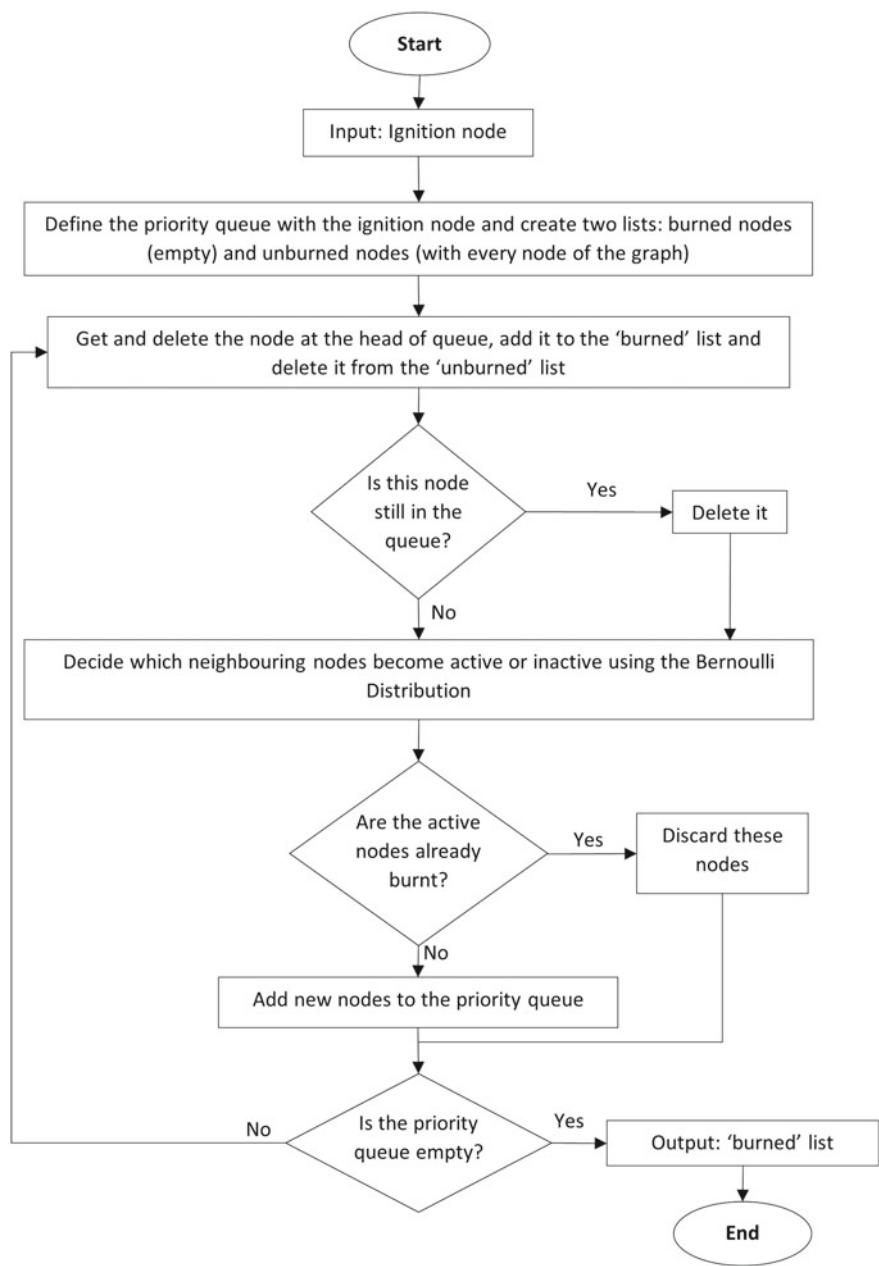
The initial step is to remove the node with the shortest propagation time from the priority queue, add it to the list of nodes that have already been visited and assume this node as the current node. If this node is still in the priority queue with a longer propagation time, it is deleted from the queue. Next, it is necessary to check the neighboring nodes of the current node and make them active or inactive using a '*Bernoulli Generator*'. This tool generates random binary numbers using a Bernoulli distribution with parameter  $\text{textit{p}}$ , that produces zero with probability  $1-p$  and one with probability  $p$ . The probability  $p$  used is determined according to equation (5) and corresponds to '*weight1*' (probability of the fire spread from one node to the other), and it is calculated a priori. For each of the neighboring nodes of the current node, it is thus essential to:

1. Check the probability ('*weight1*') concerning the arc that starts from the current node towards the neighboring node.
2. Apply the '*Bernoulli Generator*' to this probability.
3. Check the numbers obtained in 2 and define the neighboring nodes as active (if connected to the current node by an arc with '1' value) or inactive (if connected to the current node by an arc with '0' value).

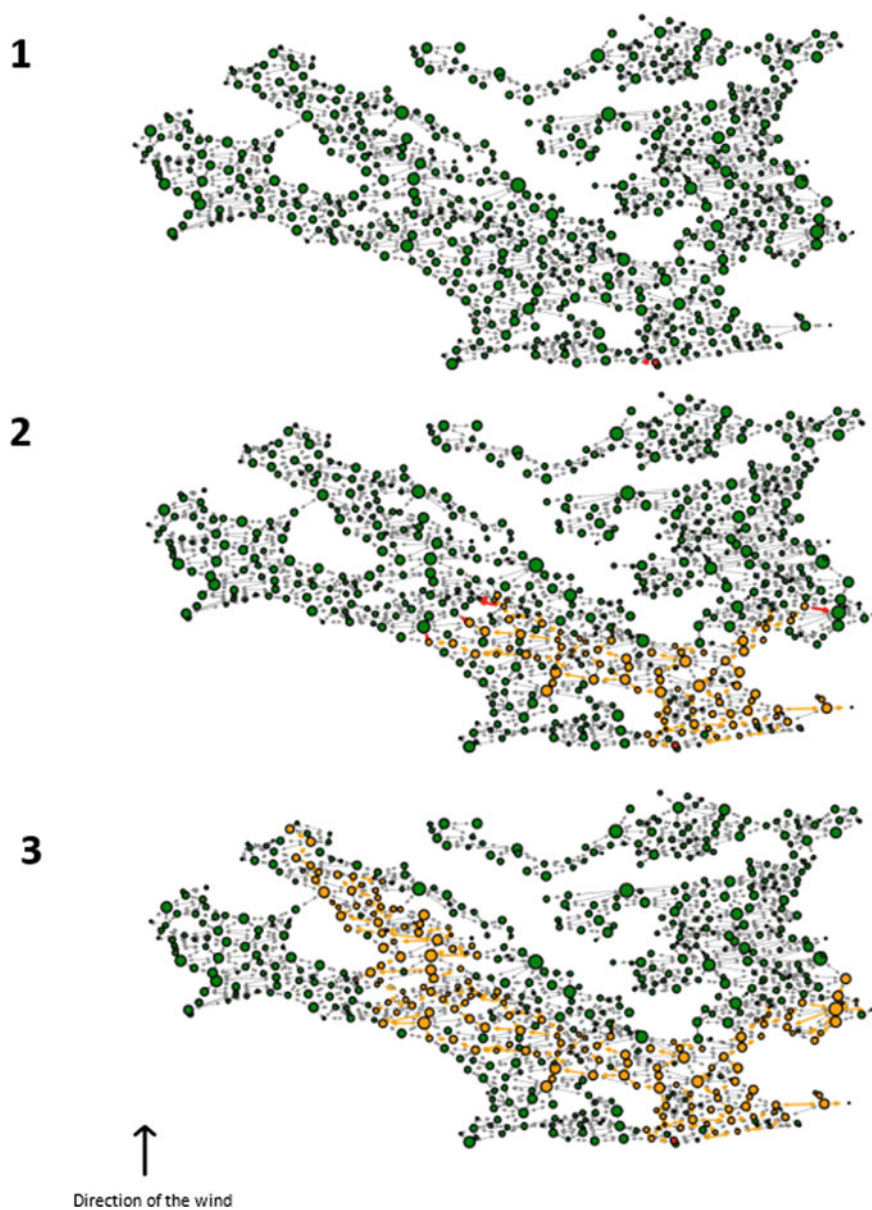
Then, it is verified, according to step 3, which of the neighboring nodes are active. If the nodes that result from this step have already been visited, they are ignored. Otherwise, they remain active and, therefore, are added to the priority queue. The propagation time associated to each node in the queue is the sum of the previous node's time with the current new one ('*weight2*'). As long as the queue is not empty, the process is repeated. Note that with this approach, there can be isolated nodes eventually, i.e., that will never become active by the '*Bernoulli Generator*' and therefore the fire does not reach them. Figure 3 shows the flowchart corresponding to the discrete event simulation model used, making it easier to understand the algorithm. The DES model was implemented using Python and the output of its application can be visualized in the actual graph where the set of all paths can be identified.

### 3 Results

The DES method was tested with different wind directions using the data of the Entre Douro e Sousa ZIF. Applying the procedure to the graph shown in Fig. 1(b) gives a real insight into how the model can work, since the distances between stands and the area of each stand are real. The total area covered by the graph is approximately 6



**Fig. 3** Flowchart of the discrete event simulation model



**Fig. 4** Entre Douro e Sousa ZIF graph with DES applied: (1) the fire started but has not spread yet; (2) the fire spreading and (3) the fire extinguished

611.85 ha. The experiment described next assumes a fixed ignition point located in the south of the forest.

The fire spread in the forest with the wind direction to the north (winds come from the south and are directed north) can be seen in Fig. 4. It is noticeable that the direction of the fire spread effectively follows the wind direction. 30.65% of the stands in the forest were burnt, which corresponds to a burnt area of approximately 2 175.89 ha (32.9% of the forest area). Wind speed was 4.67 km/h and the speed of fire spread (given by the Rothermel model) was 0.12 km/h, which is a very slow speed, consequently it is expected that the fire spread in the forest will take a long time. In fact, the fire would take approximately 511 days to burn 2 175.89 ha with a weak wind speed. Clearly, under these conditions the result is unrealistic since there are firefighting means to extinguish the fire.

For strong winds the time of propagation in the forest is expected to decrease, since the speed of fire propagation increases. Thus, for the same conditions, but with a higher wind speed (gale), there is actually a shorter time: (i) the burnt area was 2 192.01 ha which means that 33.15% of the forest was burnt; (ii) the wind speed was 64.80 km/h and the fire spread speed was 1.39 km/h; and (iii) the time of fire spread in the forest was 42.5 days. Still, 42.5 days to burn an area of 2 000 ha is a long time, nonetheless there is an increase in the speed of fire spread, as expected. Further on, this unrealistic fire spread time will be addressed by applying another method (10% of the wind speed) to calculate the speed of fire spread.

Next, in order to gauge more realistic fire spread times, the 10% wind speed rule of thumb for the fire spread speed was applied, using the same conditions. Assuming a light wind a fire spread similar to the previous one is achieved. The wind speed was 8.54 km/h and the fire spread speed was 0.85 km/h. It can be seen that using this rule, instead of Rothermel model, the fire propagation is faster. The burnt area in this situation is 1 864.6 ha which corresponds to 28.2% of the forest and the fire propagation time was 67.3 days. Applying the same conditions but for strong wind, they clearly show a faster spread. With a wind speed of 75.1 km/h and a fire spread speed of 7.51 km/h, the burnt area was 2 029.74 ha (30.7% of the forest). The fire propagation time was almost 8 days. From an initial perspective, using the 10% wind speed rule of thumb appears to achieve more realistic results. In fact, the parameters used in the Rothermel model were constant, except for the wind speed, so the characteristics of this area may not match the values used for those parameters, which might explain these differences in results.

In order to make a more in-depth analysis of the results, the model that simulates the spread of fire in the Entre Douro e Sousa ZIF network was run a thousand times with all the previous assumptions, assuming the two alternatives for the calculation of the fire spread rate: the Rothermel model and the 10% wind speed rule of thumb. The results of the 95% confidence intervals (CI) for the means for the percentage of area burnt and the total spread time can be seen in Table 2. Assuming that the wind direction is always north, it can be noticed that the percentage of burnt area practically does not change, whether the rate of fire spread is determined using the Rothermel model or the 10% wind speed rule. It also does not change when the wind speed is stronger or weaker, because the way the fire spreads depends mostly on the

**Table 2** Results obtained for the different scenarios of the wind speed regarding the time of fire spread in the ZIF graph and the percentage of burnt area

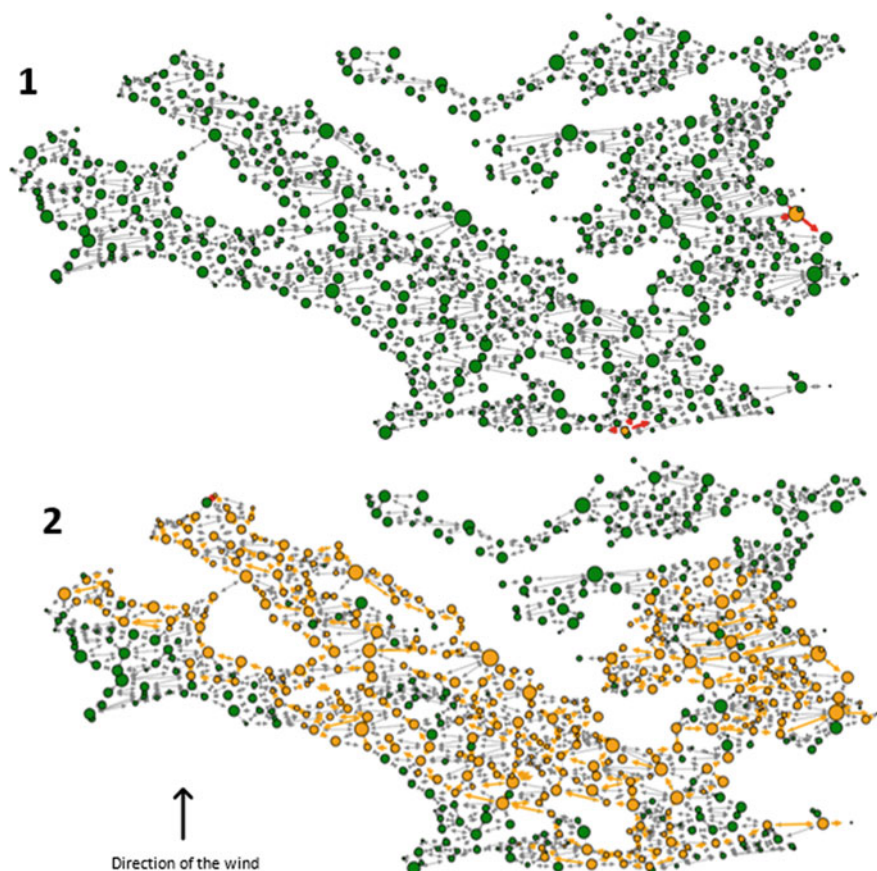
| Wind speed                          | Rothermel      |              | 10% wind speed rule |              |
|-------------------------------------|----------------|--------------|---------------------|--------------|
|                                     | Light breeze   | Gale         | Light breeze        | Gale         |
| 95% CI for percentage of burnt area | [31.1, 33.5]   | [31.1, 33.4] | [31.2, 33.4]        | [30.0, 32.2] |
| 95% CI for total time               | [253.0, 264.2] | [40.0, 41.8] | [75.0, 78.1]        | [7.3, 7.6]   |

probabilities of fire spread (*weight1*) that depend only on the wind direction that, in this case, was always the same (north). Regarding the total time of fire spread, there are noticeable differences, as it happened when the simulation was run only once, namely, when considering the 10% wind speed rule, there is a decrease in the time of propagation, both for strong and weak winds.

The units expressed for the total time of fire propagation are days. Therefore, for strong winds, and considering the speed of fire spread calculated by the Rothermel model, it is verified that the propagation time, on average, is between 40.0 to 41.8 days. Whereas, if the speed of fire spread is determined using 10% of the wind speed, this interval decreases to 7.3 to 7.6 days. Effectively, it is verified that the speed of fire spread is much lower for the Rothermel model and that can be explained by the disregard of the forest aspects. However, to burn 2 000 ha (30% of the total area) in a real situation would not take more than 40 days, so using the speed of fire spread as 10% of wind speed leads to more realistic results.

All the experiments done previously had only one ignition point and this node was always chosen strategically, that is, considering the wind direction, an ignition point was chosen that favoured the propagation of the fire in the forest. However, it is possible for a fire to have more than one ignition point and, as stated earlier, there are always areas where it is more likely for a fire to start. Next, it will be assumed there are two ignition points strategically placed, one to the south of the forest (the same node used previously) and the other to the east, and strong winds directed north (Fig. 5 (1)). The rate of fire spread considered is 10% of the wind speed. The fire spread for the two ignition points can be seen in Fig. 5 (2), where a clear connection of the two fires to a larger one can be noticed, however there may be cases where both fires do not merge. The total area burnt was 3 884 ha, which corresponds to approximately 58% of the forest and the total time of propagation was close to 10 days.

In order to evaluate the influence of multiple ignition points, the results with one ignition point (in a location prone to fire spread—south and another in a location unfavourable to fire spread - north) with those of two ignition points (south and east) are compared. To better analyse the results the model was run a thousand times for strong and weak winds, always with the same ignition points, and the results can be seen in Table 3. When applying the Z-test to the two samples of percentages of burnt



**Fig. 5** Entre Douro e Sousa ZIF graph with DES applied: (1) the fire started but has not yet spread and (2) the fire extinguished

area for each of the three scenarios, a non-significant p-value ( $> 0.05$ ) was obtained, indicating that there are no significant differences regarding the average percentage of burnt area, for light breeze versus gale.

The results presented in Table 3 for a fixed ignition point located south of the forest are the same as in Table 2, for the 10% wind speed rule. There is a slight decrease in the percentage of burnt area in the case of the two ignition points, when compared to one ignition point located in the south, and a strong decrease in the time of fire spread. There are two potential explanations for why the fire spread time decreased in the scenario of the two ignition points: (i) the decrease in the percentage of burnt area, since the fire propagation time decreases with the decrease in burnt area, and (ii) the existence of the two ignition points which causes the fire to propagate at the same time in two different places of the forest thus decreasing its propagation time.

**Table 3** Results obtained for the Entre Douro e Sousa ZIF graph with one and two ignition points for two different wind speeds

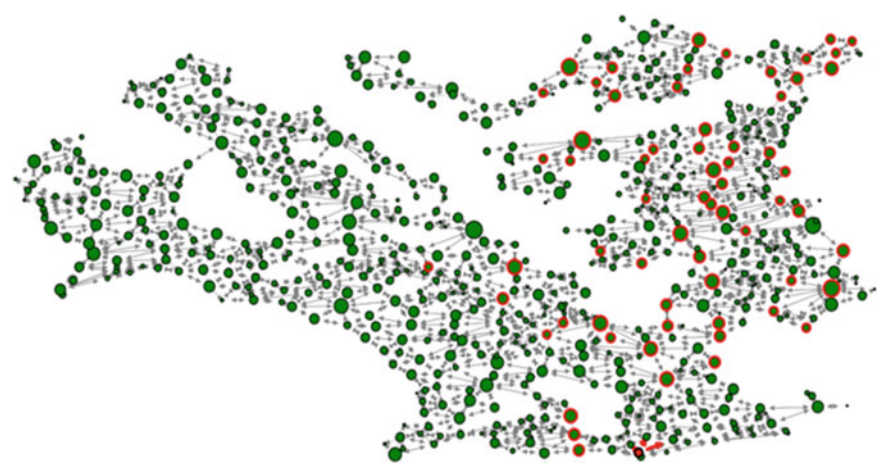
|                                     | Fixed ignition point (south of the forest)               |              |                |
|-------------------------------------|----------------------------------------------------------|--------------|----------------|
|                                     | Light breeze                                             | Gale         | Z test p-value |
| 95% CI for percentage of burnt area | [31.2, 33.4]                                             | [30.0, 32.2] | 0.147          |
| 95% CI for total time               | [75.0, 78.1]                                             | [7.3, 7.6]   | < 0.005        |
|                                     | Fixed ignition point (north of the forest)               |              |                |
|                                     | Light breeze                                             | Gale         | Z test p-value |
| 95% CI for percentage of burnt area | [5.5, 5.9]                                               | [5.4, 5.8]   | 0.315          |
| 95% CI for total time               | [26.3, 27.7]                                             | [2.9, 3.1]   | < 0.005        |
|                                     | Two fixed ignition points (south and east of the forest) |              |                |
|                                     | Light breeze                                             | Gale         | Z test p-value |
| 95% CI for percentage of burnt area | [29.7, 31.6]                                             | [28.9, 30.8] | 0.231          |
| 95% CI for total time               | [45.3, 47.4]                                             | [5.3, 5.6]   | < 0.005        |

The results were in line with what was expected since with two ignitions points of fire it would be expected that the fire spread is faster. Even so, the percentage of burnt area decreases slightly, contrary to what was expected.

The results presented are, as it was previously mentioned, for two fixed ignition points and one thousand runs, i.e., as a new network is formed, the ignition points remain unchanged, one to the south and the other to the east of the forest. The ignition point located to the south of the forest is favourable to fire spread since the wind direction is north, so there is a larger area through which the fire can effectively propagate. For the northern ignition point the opposite happens, making it unfavourable to the fire spread, which is corroborated by the low percentages of burnt area.

According to the literature ([11, 16, 17]), the existence of roads, the terrain elevation and the area, are some of the influential parameters in fire ignition and there is information regarding these three parameters for each of the stands. In order to visualise these three parameters, it was decided to assume an elevation above 290m and an area greater than 9 ha. Therefore, in Fig. 6, the stands with these conditions, and with roads, are represented with a red border. It was considered that a stand without those conditions (i.e. without a red border) has a zero probability of ignition, while those identified with a red border have an equal probability of being selected as a point of ignition.

Since the elevation of the forest is higher to the east, the risk of fire ignition is also higher to the east where most of the stands have roads, higher elevation and larger area. Now, in order to analyse the results, the model was run a thousand times with variable ignition point within those that presented higher ignition risk, that is, in each



**Fig. 6** Entre Douro e Sousa ZIF graph: fire ignition risk map

**Table 4** Results obtained for the Entre Douro e Sousa ZIF graph with one variable ignition point for two different wind speeds and taking into account only the ignition points with higher risk

| Wind direction: North               | Light breeze | Gale        | Z test p-value |
|-------------------------------------|--------------|-------------|----------------|
| 95% CI for percentage of burnt area | [9.4, 10.7]  | [9.1, 10.4] | 0.486          |
| 95% CI for total time               | [21.7, 23.4] | [2.8, 3.1]  | < 0.05         |

run the ignition point changes to any node identified in Fig. 6. It was considered a wind direction to north and the results of these experiments can be seen in Table 4.

When comparing a fixed ignition point in the south of the forest (Table 3) with a variable ignition point (Table 4) there is a clear decrease in the percentage of burnt area as well as in the fire propagation time when the ignition points become variable, as expected since the points may not be favourable to the fire propagation, considering the direction of the wind.

## 4 Conclusions

The work described in this paper focused on the development of a Discrete Event Simulation method to model fire propagation in a graph, which proved to be effective and adequate. The model was developed in order to obtain a greater insight into how fire can spread through a specific landscape, and how the factors considered can affect the behaviour of the fire and the time it takes to burn the area. However, there are some limitations that can be pointed out. This method is greatly influenced by the probabilities of fire spread and therefore these should be defined as accurately as

possible, but it should be noted that to obtain these values in a real situation is a complex process. The wind was chosen, as it is the parameter that most influences these probabilities, but as no other aspect of the forest was considered, the characterisation of the area should be improved.

Two other points that should be considered in further studies are the way the probabilities of fire ignition are defined and how to include the volatility of the wind in the model. The former was obtained in a very simple way, which, in future work, would require improvement since assigning a probability of zero to a certain number of nodes may be a too drastic approach. In the case of the wind, it was considered constant (both wind speed and direction) throughout the forest and during the entire time the fire was spreading. Clearly this assumption does not correspond to reality, so assigning greater volatility to the wind should be included in further analyses, based on, preferably, real data regarding wind behaviour in that area. If real data is used, the impact of wind speed and direction on the speed of fire spread can be included in the model in a more realistic manner, thus leading to more accurate results.

Furthermore, for future work fire projections could be incorporated, i.e., for certain conditions the fire can be projected in such a way that a new ignition in another place of the forest occurs, as well as the impact of fire fighting resources in the propagation of the fire, since it is expected that a forest fire is not left unattended until it extinguishes by itself.

The data obtained from the Paiva and Entre Douro e Sousa ZIF were very important for the verification and incorporation of the study in real situations. It is also important to use these data to improve the model (e.g., the calculation of probabilities) and to study the impact of other factors (e.g., slope, vegetation). The models applied showed results close to what would happen in reality, evidencing the importance and suitability of the study. Wind speed and direction, as well as the location of the fire ignition point, are factors that greatly affect fire behaviour, and this was verified in this study.

Forest fires unfortunately affect the planet more and more, therefore, their study and understanding are of great importance in order to help to reduce or eliminate them. The model applied should thus contribute to the understanding of this issue, as a problem that highly impacts the environment and the life of the communities.

**Acknowledgements** This work was supported by the Portuguese National Funding Agency for Science, Research and Technology (FCT), within the Center for Research and Development in Mathematics and Applications (CIDMA), project UIDB/04106/2020; and through project PCIF/GRF/0141/2019-An Optimization Framework to reduce Forest Fire.

## References

1. Moritz, M.A., Batllori, E., Bradstock, R.A., Gill, A.M., Handmer, J., Hessburg, P.F., Leonard, J., McCaffrey, S., Odion, D.C., Schoennagel, T., Syphard, A.D.: Learning to coexist with wildfire. *Nat.* **515**(7525), 58–66 (2014)

2. Hajian, M., Melachrinoudis, E., Kubat, P.: Modeling wildfire propagation with the stochastic shortest path: a fast simulation approach. *Environ. Model. & Softw.* **73**–88 (2016)
3. Andrews, P.L.: The Rothermel surface fire spread model and associated developments: a comprehensive explanation. In: *Gen. Tech. Rep. RMRS-GTR-371*. Forte Collins, CO:U.S. Department of Agriculture, Forest Service, Rocky Mountain Research Station, p. 121 (2018)
4. Cruz, M.G., Alexander, M.E.: The 10% wind speed rule of thumb for estimating a wildfire's forward rate of spread in forests and shrublands. *Ann. For. Sci.* **76**, n° 44 (2019)
5. Mutthulakshmi, K., Wee, M.R.E., Wong, Y.C.K., Lai, J.W., Koh, J.M., Acharya, U.R., Cheong, K. H.: Simulating forest fire spread and fire-fighting using cellular automata. *Chin. J. Phys.* (2020)
6. Alexandridis, A., Vakalis, D., Siettos, C.I., Bafas, G.V.: A cellular automata model for forest fire spread prediction: the case of the wildfire that swept through Spetses Island in 1990. *Appl. Math. Comput.* 191–201 (2008)
7. Islam, M.R., Saidur, R., Rahim, N.A.: Assessment of wind energy potentiality at Kudat and Labuan, Malaysia using Weibull distribution function. *Energy* **36**, n° 2, 985–992 (2011)
8. Seguro, J.V., Lambert, T.W.: Modern estimation of the parameters of the Weibull wind speed distribution for wind energy analysis. *J. Wind. Eng. Ind. Aerodyn.* **85**(1), 75–84 (2000)
9. Delmar-Morgan, E.L.: The Beaufort scale. *J. Navig.* **12**(01), 100–102 (1959)
10. Bar Massada, A., Syphard, A.D., Stewart, S.I., Radeloff, V.C.: Wildfire ignition-distribution modelling: a comparative study in the Huron? Manistee National Forest, Michigan, USA. *Int. J. Wildland Fire* **22**(2), 174–183 (2013)
11. Badia, A., Serra, P., Modugno, S.: Identifying dynamics of fire ignition probabilities in two representative Mediterranean wildland-urban interface areas. *Appl. Geogr.* **31**, 930–940 (2011)
12. Vásquez-Serrano, J.I., Peimbert-García, R.E., Cárdenas-Barrón, L.E.: Discrete-event simulation modeling in healthcare: a comprehensive review. *Int. J. Environ. Res. Public Health* (2021)
13. Hill, R.R., Miller, J.O., McIntyre, G.A.: Applications of discrete event simulation modeling to military problems. In: *Proceedings of the 2001 Winter Simulation Conference* (2001)
14. Turner, C.J., Hutabarat, W., Oyekan, J., Tiwari, A.: Discrete event simulation and virtual reality use in industry: new opportunities and future trends. *IEEE Trans. Hum. Mach. Syst.* (2016)
15. Filippi, J.B., Morandini, F., Balbi, J.H., Hil, D.: Discrete event front tracking simulator of a physical fire spread model. *SIMULATION*, SAGE Publications **86**, n° 10, 629–644 (2009)
16. Rodríguez-Martínez, A., Vitoriano, B.: Probability-based wildfire risk measure for decision-making. *Math.* **557** (2020)
17. Catry, F.X., Rego, F.C., Bação, F., Moreira, F.: Modeling and mapping wildfire ignition risk in Portugal. *Int. J. Wildland Fire* 921–931 (2009)
18. A. F. d. V. d. Sousa, “afvs” [Online]. <http://www.afvs.ws/index.html>. Accessed 15 Dec 2022

# Multiobjective Evolutionary Clustering to Enhance Fault Detection in a PV System



Luciana Yamada, Priscila Rampazzo, Felipe Yamada, Luís Guimarães, Armando Leitão, and Flávia Barbosa

**Abstract** Data clustering combined with multiobjective optimization has become attractive when the structure and the number of clusters in a dataset are unknown. Data clustering is the main task of exploratory data mining and a standard statistical data analysis technique used in many fields, including machine learning, pattern recognition, image analysis, information retrieval, and bioinformatics. This project analyzes data to extract possible failure patterns in Solar Photovoltaic (PV) Panels. When managing PV Panels, preventive maintenance procedures focus on identifying and monitoring potential equipment problems. Failure patterns such as soiling, shading, and equipment damage can disturb the PV system from operating efficiently. We propose a multiobjective evolutionary algorithm that uses different distance functions to explore the conflicts between different perspectives of the problem. By the end, we obtain a non-dominated set, where each solution carries out information about a possible clustering structure. After that, we pursue a-posteriori analysis to exploit the knowledge of non-dominated solutions and enhance the fault detection process of PV panels.

**Keywords** Multiobjective · Clustering · Photovoltaic systems · Fault detection

## 1 Introduction

In the last years, the Paris Agreement has defined the necessary targets to limit global warming using renewable energy. Despite decreasing solar panels' costs, the panels' efficiency is still reduced compared with other energy sources [1]. The industry has been working to improve the overall performance of photovoltaic (PV) systems. However, issues still need to be addressed concerning reliability, unforeseen outages,

---

L. Yamada (✉) · F. Yamada · L. Guimarães · A. Leitão · F. Barbosa  
INESC TEC, Faculdade de Engenharia, Universidade do Porto, Porto, Portugal  
e-mail: [luciana.o.yamada@inesctec.pt](mailto:luciana.o.yamada@inesctec.pt)

P. Rampazzo  
Faculdade de Ciências Aplicadas, Universidade Estadual de Campinas, Limeira, Brazil

and high operation and maintenance (O&M) costs, hindering a lean integration in the electrical grid.

A PV system comprises one or more solar panels, connected in series or parallel, combined with an inverter, and a utility grid, among other components. PV panels convert sunlight into electrical energy, and faults may affect performance. A PV fault decreases the system's performance (reduces the output) or interrupts the system's operation. Faults can occur by different factors such as errors in the Maximum Power Point Tracking, shadowing, degradation, and electrical disconnection. Such failures can reduce the instantaneous power generated by the PV plant or permanently degrade the overall asset. In this sense, implementing advanced analytical tools to diagnose failures is crucial in guaranteeing the asset's reliability and performance. In addition, detecting and diagnosing potential failures is also decisive to reduce the costs associated with O&M and the system unavailability.

Although different techniques have been reported in the literature, the fault detection process must deal with some problems [2]. Regarding data availability, accurately representing the electrical characteristics under different fault conditions is challenging, and most fault diagnosis models require the precise division of the fault samples. Fault detection methods based on classification models can effectively reveal and categorize faults. However, specialized human knowledge or complex and costly equipment are needed to establish the diagnostic model and fault samples. Effective clustering of fault samples is a prerequisite for establishing an efficient detection model, but the scarcity of fault samples in real operational data makes this task difficult [3].

Considering as challenges of the fault detection problem in PV systems: (i) the difficulty of categorizing faults and having data labeled according to their different types, and (ii) the lack of information regarding the distribution of data in high-dimension, this paper proposed a Multiobjective Evolutionary Clustering Algorithm (MECA) to detect faults in PV systems through the following approach:

- i Application of a multiobjective algorithm for the clustering task—to capture the different clustering perspectives through different objective functions to overcome the lack of knowledge about the data distribution format due to the diversity of failures and high-dimensional spaces.
- ii Proposal of a-posteriori analysis—associating data with known labels to the results of the multiobjective algorithm for the clustering task; the aim is to obtain a cluster-sharing matrix to enhance the task of attributing labels to initially unlabeled data.
- iii Validation of the method—use the multiobjective approach for clustering and enhance the assignment of missing labels in databases with few known labels—through the measurement of accuracy and confusion matrix.

The main contribution of this paper is the proposal of a-posteriori analysis that demonstrated to be able to label data and explore complementary information for data clustering.

This paper is organized as follows. Section 2 gives a description of fault detection techniques. Section 3 presents the Multiobjective Evolutionary Clustering algorithm. Section 4 reports the experimental results. Section 5 summarizes the conclusion of this work and future research.

## 2 Fault Detection

According to [4], fault detection techniques are divided into visual and thermal, and electrical methods. Visual and thermal imaging methods help in situations where faults are challenging to detect in the visual inspection process. These methods analyze images collected by cameras that can detect temperature differences in the PV module and recognize the failures' exact location. The electrical methods use the measurements of the electrical output such as the current and voltage of the PV system, to diagnose the faults. Among the electrical methods, some of the techniques found are statistical and signal processing approaches, current-voltage characteristics analysis, power losses analysis, and artificial intelligence (AI) techniques [5]. In this work, we will focus on AI techniques.

The AI techniques has shown significant advances and contributions to several scientific areas. One of the main uses of fault detection is related to supervised methods. In [6], a Monitoring System (MS) is presented to measure the electrical and environmental variables to produce instantaneous and historical data. Integrated with the MS, an Auto-Regressive with an Exogenous input model is used to detect faults in the system, such as short-circuit, open-circuit, partial shadowing, and degradation. Regarding the classification of the faults, different supervised models were compared.

Unsupervised methods, such as data clustering, are applied to group data with similar characteristics, differentiating data with failures and data in regular operation. In [7], the K-means algorithm is used to cluster thermal images to detect and localize damage. Elbow method and mean silhouette method are used to define the best number of clusters. K-means proved to detect faults in images and can be integrated into the thermal drone system. In [8], an unsupervised method, density-based spatial clustering of applications with noise algorithm, is used for clustering faults in a PV system. The normalization proposed in [9] is used to avoid overlap of the clusters. The approach identifies the faults of open-circuit and short-circuits.

In addition to the mentioned methods, semi-supervised approaches are applied in the literature. The main objective of these algorithms is to combine many labeled samples with a few unlabeled ones to develop more efficient models [10]. These algorithms try to improve the performance of supervised or unsupervised learning using a combination of methods and information generated by each other [11]. In [9], a Graph-Based Semi-supervised Learning (GBSSL) method is proposed for fault detection and classification in PV arrays. The normalization based on the panels' standard test condition was proposed and proved to help avoiding overlap clusters under normal and fault conditions. The GBSSL demonstrated that the faults could

be correctly detected during weather changes or PV arrays degradation. In this work, the authors consider the short and open-circuit faults.

The Fuzzy C-Means (FCM) [12] algorithm is a clustering algorithm based on the fuzzy division of an objective function that groups samples of failures by the degree of membership through the Euclidean distance between the sampling point and the center of the corresponding cluster, which makes the algorithm limited to processing datasets with spherical shapes. The authors of [3] propose the function of the Gaussian kernel (GK) in the FCM algorithm to map data in the high-dimensional space and transform the non-linear information of the original space into a linear problem, indicating a significant improvement in the applicability and the algorithm clustering accuracy. The GK-FCM algorithm performs unsupervised clustering and labeling fault samples under typical fault conditions. The labeled fault samples serve as input to a probabilistic neural network based model to facilitate intelligent diagnosis of PV array faults.

Although works with similar intentions have been proposed in the literature, the approach presented here is new: treating the fault detection PV system by analyzing the results of a multiobjective clustering algorithm. Unlike the existing works, we extract additional information from clustering in a-posteriori analysis, considering a few labeled samples, and use this information to classify the failure types for unlabeled data.

### 3 Methodology

In this section, an overview of MECA is presented. Based on the algorithm presented in [13], we propose an algorithm to cluster PV system data and find possible labels to unlabeled data in a-posteriori analysis. Our algorithm has some changes from the original algorithm regarding the initialization and the mutation operators. The general framework of MECA is outlined in Algorithm 1.

The individual of MECA represents a way of grouping the status of a PV system for a specific irradiance range. Each status can be associated with one type of fault or a normal condition. Status with similar characteristics should be clustered together. The individual was coded using a locus-based adjacency coding. The methods used in initialization introduce individuals with characteristics explored in the two objective functions: Compactness and Connectivity. K-means aims to minimize the sum of squared errors of each sample to the nearest centroid. The centroids represent the average of the samples in each cluster. Therefore, this algorithm highlights structures of compact clusters. Similarly, K-medoids prioritize compact clusters with less sensitivity to outliers. Kruskal Algorithm introduces individuals with smaller distances between the samples, highlighting the connectivity of the data.

Next, we use the NSGA-II [14] to sort and select individuals from the population in the binary tournament to compose the parents population  $P_g$ . The selected set of parents go through crossover and mutation to create an offspring population  $Q_g$ . The uniform crossover was used to combine genes from two parents. The mutation oper-

ator proposed in this work emphasizes the characteristics of the objective functions. Relying on the data distribution, one of the objective functions or a mix of the two can better represent the data. For example, K-means and K-medoids can not create a proper individual due to the spherical characteristic in data with the elongated or connected format. In contrast, Kruskal Algorithm prioritizes connectivity, and it is hard to deal with data with different types of distribution or even in compact or spherical shapes. Therefore, the neighborhood mutation was developed to change the assignment of a sample to the cluster of its nearest neighbor, exploring the connection between samples. Similarly, the centroid mutation explores the characteristic of compact clusters, changing a sample to another closest centroid. These designed mutation operators help maintain the genetic diversity of the population. In our fault detection problem, as we can not know the data distribution and consider different irradiance ranges, the proposed objective function and mutation operators is a proper choice to capture distinct patterns.

The combined population  $R_g = P_{g-1} \cup Q_g$ , is sorted according to non-dominated classification and crowding distance to choose exactly  $TP$  population members to the next population  $P_g$  [14].

After obtaining the non-dominated frontier, a-posteriori analysis is applied to exploit the knowledge of the solutions and enhance the fault detection process of PV panels. For this purpose, we combine the information obtained from the non-dominated solutions with a small percentage of labeled samples.

More details of the implementations of the coding, evaluation, operators and a-posteriori analysis are described in the following.

---

**Algorithm 1: MECA**


---

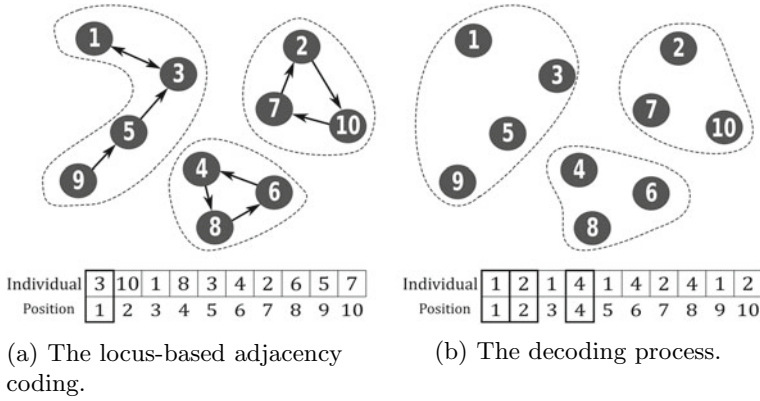
```

1 Initialize population  $P_0$  with size  $TP = Q_{Kmeans} + Q_{Kmedoids} + Q_{MST}$ 
2 for  $g = 1, \dots, G$  do
3   Select the sets of parents,  $p_1$  and  $p_2$ , through binary tournament
4   Create  $F_c$  offsprings from  $p_1$  and  $p_2$  (uniform crossover)
5   Create  $F_m$  offsprings from  $p_1$  and  $p_2$  (centroid and neighborhood mutation)
6    $Q_g = F_c \cup F_m$  Create  $R_g = P_{g-1} \cup Q_g$ 
7   Evaluate  $R_g$  in Compactness and Connectivity
8   Select  $TP$  individuals of  $R_g$  through rank and crowding distance to form the next
   population  $P_g$ 
9 end
10 Return the non-dominated solutions

```

---

All the parameters were defined experimentally. The population size ( $TP$ ) was set to 100, and the number of generations ( $G$ ) to 100. The initial population is generated using the following K-means [15], K-medoids [16], and the Kruskal Algorithm [17]. We define the number of solutions of each method, Kruskal ( $Q_{MST}$ ), K-means ( $Q_{Kmeans}$ ), and K-medoids ( $Q_{Kmedoids}$ ), respectively to 24, 38 and 38. For each method, an individual will be created considering the chosen range of the cluster where  $[K_{min}, K_{max}] = [2, 10]$ .



**Fig. 1** Illustration of the locus-based adjacency coding and the decoding process

## Coding

In our dataset, each observation corresponds to a status of a PV system, measured by attributes related to weather conditions and electrical variables. The status indicates whether the system operates in a normal condition or has any faults. Thus, each individual in the population corresponds to a way of grouping the  $N$  observations of the dataset. To represent the individual, we used locus-based adjacency coding [13], where an individual is a vector  $[a_1, a_2, \dots, a_N]$ , where  $a_i$  is associated with sample  $i$ -th and represents one of the nodes of the graph. The value of an element  $a_i = n$  indicates that sample  $i$  is linked to sample  $n$ . These links create the node's connection and, in the end, generate the final clusters in the decoding process. With this representation, the number of clusters in individuals can be flexible during the search process. In Fig. 1a, we represent an individual with locus-based adjacency coding. For example, position one is associated with node and sample 1. The value 3 indicates that sample 1 has a link with sample 3, as we can see in the graph. In Fig. 1b, the individual is decoded using one of the samples as the root (smallest value) to group samples assigned to the same cluster. In the decoding process, the roots are 1, 2 and 4.

## Solution Evaluation

To evaluate each individual ( $k = 1, \dots, TP$ ) of the population and express different characteristics of the data, we use two objective functions based on the MOCK [13]: Compactness and Connectivity. As an objective reflecting the cluster compactness (1), we minimize the Euclidean distance  $\delta(i, \mu_j)$  between each sample  $i$  of the cluster  $C_j$  with its centroid  $\mu_j$ :

$$\text{Compactness} = \sum_{C_j \in C} \sum_{i \in C_j} \delta(i, \mu_j). \quad (1)$$

To represent the cluster connectivity (2), we apply a penalty factor that considers the distance from the nearest neighbors. The nearest neighbors are calculated using the euclidean distance between data points. In the objective function,  $nn_{i,j}$  is the  $j$ -th nearest neighbor of a sample  $i$ ,  $L$  is the number of neighbors that are considered to compute the metric, and  $N$  is the size of the samples in the dataset. For the parameter  $L$ , we set it to 10, as suggested in [13]. This objective function should be minimized.

$$\text{Connectivity} = \sum_{i=1}^N \sum_{j=1}^L p_{i,nn_{i,j}}. \quad (2)$$

The penalty factor use the following rules:

$$p_{i,nn_{i,j}} = \begin{cases} \frac{1}{j}, & \text{if } nn_{i,j} \text{ does not belong to the same cluster of the sample } i \\ 0, & \text{otherwise.} \end{cases}$$

## Non-dominated Sorting Genetic Algorithm II (NSGA-II)

As defined in [14], in NSGA-II each individual ( $k = 1, \dots, TP$ ) is associated with two attributes:  $rank_k$  and  $distance_k$  (non-dominated classification and crowding distance). If two solutions have different non-domination levels (different non-dominated frontiers), we choose the solution  $k$  with the lower  $rank_k$ . Otherwise, if two  $k_1$  and  $k_2$  solutions belong to the same frontier ( $rank_{k_1} = rank_{k_2}$ ), then we prefer the solution that is located in a less crowded region (that is, higher  $distance_k$ ) [18]. The crowding distance represents the sum of the normalized distances of the nearest neighbors of the solution along each objective. Large distance values are assigned for the extreme solutions, and the crowding distance is calculated for the rest. Formally, given a solution  $m$ , the distance  $d_m$  is defined as:

$$d_m = \sum_{f=1}^F \frac{g_f(m+1) - g_f(m-1)}{g_f^{max} - g_f^{min}} \quad (3)$$

Where  $F$  is the number of objective functions,  $m+1$  and  $m-1$  are the nearest neighbors of a solution  $m$ ,  $g_f^{max}$  and  $g_f^{min}$  the maximum and minimum values of each function  $g_f$ , considering individuals from the same frontier (same  $rank$ ).

Non-dominated classification and crowding distance are used in the binary tournament selection algorithm to create the parents' population and to select the population for the next generations.

## Crossover and Mutation

A population of parents is built through a binary tournament. Pairs of parents from this population are selected to generate pairs of offspring. To compose the offspring population, we define the solutions generated by the crossover ( $F_c$ ) to 20 and by the mutation ( $F_m$ ) to 20. The uniform crossover is applied to each pair of selected solutions, thus generating two offspring. The binary crossover mask with a uniform distribution is created; each individual has a 50% chance to copy the gene from one and 50% from the other.

The mutation process complements the crossover, allowing a more extensive search space to be explored. In the literature, mutation, and crossover operators for the clustering problem are presented in [19]. The mutation is based on the nearest neighbors in the algorithms that use locus-based adjacency graph encoding. In the present work, two types of mutation are proposed to emphasize the characteristics of objective functions: centroid mutation and neighborhood mutation.

The centroid mutation aims to assign the selected sample to the cluster with the nearest centroid. For each offspring:

1. Select randomly one sample  $i$  of the offspring.
2. Calculate the Euclidean distance between sample  $i$  and all the other centroids of the offspring.
3. Assign the sample  $i$  to a new cluster, considering the nearest centroid.

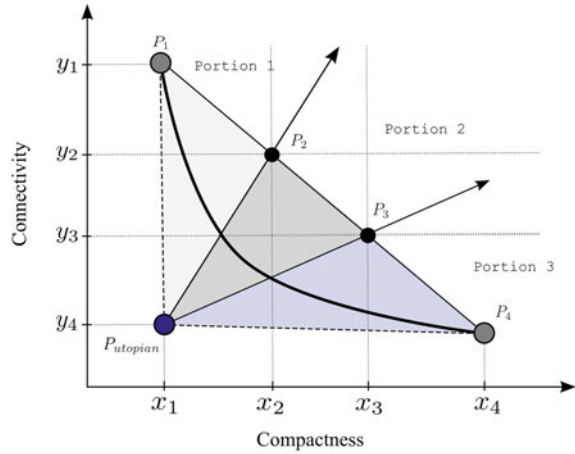
The neighborhood mutation aims to assign the selected sample to the cluster of the nearest neighbor. For each offspring:

1. Select randomly one sample  $i$  of the offspring.
2. Identify the  $L$  nearest neighbor of the sample  $i$  through the k-nearest neighbors algorithm.
3. Assign the sample  $i$  to a new cluster, considering the nearest neighbor.

## A-posteriori analysis

Following the MECA's execution, a-posteriori analysis examines the non-dominated solutions to gather insights into the clustering process. This analysis aims to determine whether we can leverage the acquired information to identify potential labels for unlabeled data. In the first step of a-posteriori analysis, we divide the non-dominated frontier into three portions, as illustrated in Fig. 2.

**Fig. 2** Division of the non-dominated frontier into three portions. The gray solutions represent the extremes of the frontier, and the blue solution represents the Utopian solution



From the extreme solutions of the frontier ( $P_1$  and  $P_4$ ), we find the corresponding Utopian solution ( $P_{\text{utopian}}$ ) and determine the equation of a line through the two extreme points. Then, we divide the frontier into three portions, finding lines that intersect the Utopian solution and the points  $P_2$  and  $P_3$ . This division aims to group similar data distribution in the non-dominated frontier. Portion 1 corresponds to solutions with better values in Compactness, while Portion 3 corresponds to solutions with better values in Connectivity. On the other hand, Portion 2 has solutions with a mix of characteristics of the two objective functions. MECA creates individuals with different numbers of clusters. As presented in [13], solutions with a minimum value in Compactness result in solutions with a high number of clusters, while a minimum value in Connectivity corresponds to solutions with a low number of clusters. The idea of dividing the frontier is to maintain the distributions of each portion as much as possible, not mixing the structure of a solution with a high number of clusters with another with a low number of clusters. Furthermore, as we do not know the data distribution of the fault detection problem, we want to know what portion better represents our data.

Besides the information provided by the non-dominated solutions, we used a small percentage of samples with known conditions (labeled data) to guide data labeling and extract complementary information regarding the clustering. These samples could be seen as faults and normal data identified by the specialists and can be used as a reference to help the decision-maker in new unknown cases.

For each portion of the non-dominated frontier, we construct a cluster-sharing matrix with dimension  $n \times q$ , where  $n$  is the total number of samples in the dataset and  $q$  are the labeled samples. The matrix elements  $c_{i,j}$  represent the number of times a determined sample  $i$  was allocated in the same cluster as the sample  $j$ . If two samples are frequently grouped, they can be considered similar and in the same cluster. From this cluster-sharing matrix (4), we can extract probability information as the chance of a sample being assigned in the same cluster as a sample with a

known condition. To find the probability that a sample  $n$  belongs to the same group as a sample  $q$ , we divide each element of the cluster-sharing matrix by the total sum of each row.

$$MP_{n,q} = \begin{pmatrix} \frac{c_{1,1}}{\sum c_{1,q}} & \frac{c_{1,2}}{\sum c_{1,q}} & \cdots & \frac{c_{1,1}}{\sum c_{1,q}} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{c_{n,1}}{\sum c_{n,q}} & \frac{c_{n,2}}{\sum c_{n,q}} & \cdots & \frac{c_{n,q}}{\sum c_{n,q}} \end{pmatrix}. \quad (4)$$

Accuracy (ACC) is used to measure the correct assignment of the labels for the unlabeled data. In our experiments, part of the samples are considered unlabeled. Thus, we use the actual label as a reference value to evaluate the performance of a-posteriori analysis. To calculate the ACC, we consider the number of observations a-posteriori analysis correctly assigned the labels ( $U_c$ ) to the total number of unlabeled samples ( $T_u$ ).

$$ACC = \frac{U_c}{T_u}. \quad (5)$$

## 4 Experimental Results

In this section, we will present the photovoltaic fault dataset and the computation results obtained from the experiments.

### 4.1 Photovoltaic Fault Dataset

We used a real-world dataset in these experiments, publicly available<sup>1</sup> [6]. The dataset was generated from a PV plant simulator that describes the system behavior. The data was collected and labeled, including faults (degradation, short-circuit, open-circuit, and shadowing) and normal condition. In this work, we consider the Normal data (No), Short-circuit (Sc), and Open-circuit faults (Oc). The Normal condition is associated with data without any fault. Open-circuit fault generates an interruption in the circulation of electric current due to a disconnection in the system [6]. Short-circuit fault occurs when a low impedance path appears along the system [6]. In this work, shadowing is not considered because it can cause bad data and lead to incorrect fault detection [9].

The dataset contains six attributes: irradiance, PV module temperature, and voltage and current output for both PV strings. To analyze the applicability of a multiobjective evolutionary clustering algorithm to fault detection, we divide the irradiance into six intervals between 801.116–936.475 W/m<sup>2</sup>. It is essential to mention that in

<sup>1</sup> [https://github.com/clayton-h-costa/pv\\_fault\\_dataset](https://github.com/clayton-h-costa/pv_fault_dataset).

**Table 1** The number of data points for each irradiance range. “T” = total number of labeled or unlabeled data, “No” = Normal data, “Sc” = Short-circuit faults, and “Oc” = Open-circuit faults

| Range | Irradiance (W/m <sup>2</sup> ) | Number of data points |     |     |     |           |     |     |     |
|-------|--------------------------------|-----------------------|-----|-----|-----|-----------|-----|-----|-----|
|       |                                | Labeled               |     |     |     | Unlabeled |     |     |     |
|       |                                | T                     | No  | Sc  | Oc  | T         | No  | Sc  | Oc  |
| 1     | 801.116–858.668                | 301                   | 146 | 118 | 37  | 1201      | 583 | 473 | 145 |
| 2     | 858.668–893.252                | 301                   | 202 | 15  | 84  | 1201      | 805 | 61  | 335 |
| 3     | 893.252–911.252                | 300                   | 130 | 169 | 1   | 1201      | 519 | 679 | 3   |
| 4     | 911.252–922.624                | 301                   | 101 | 196 | 4   | 1201      | 403 | 782 | 16  |
| 5     | 922.624–929.616                | 300                   | 58  | 35  | 207 | 1202      | 232 | 142 | 828 |
| 6     | 929.616–936.475                | 301                   | 57  | 22  | 222 | 1200      | 228 | 85  | 887 |

this dataset, the short-circuit and open-circuit faults occurred at a high irradiance level. In real situations, there is much more unlabeled data than labeled because the labeling process requires specific technical knowledge, is time-consuming, and can be costly [2]. A database was adapted to apply the a-posteriori analysis associated with data with known labels and evaluate the possibility of assigning labels to unlabeled data through this proposal. We randomly selected 1500 samples for each irradiance range and considered 80% of the samples unlabeled to test our data labeling approach. The labeled data is considered as the samples with known conditions, whereas the condition of the unlabeled data will be determined through a-posteriori analysis. Table 1 presents the number of data points for each irradiance range and each condition.

## 4.2 Computation Results

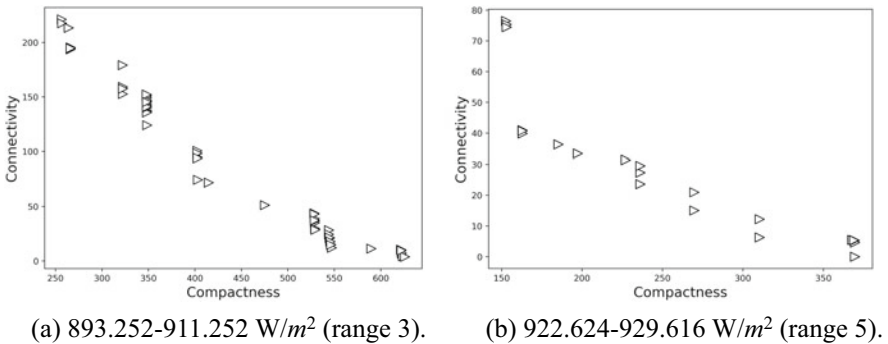
Due to the genetic algorithm’s stochastic nature, each irradiance range’s dataset is run ten times. Then, a-posteriori analysis was applied considering the non-dominated frontier information and the samples with the known conditions (labeled samples), which are records of the PV system with some fault or normal data. In a-posteriori analysis, we identify the unknown operation condition based on the probability information extracted from the cluster-sharing matrix.

In Table 2, we reported the average, the standard deviation, and the minimum and maximum obtained accuracy values of the data labeling from 10 runs based on a-posteriori analysis.

Since we had small deviations for all intervals, we selected two intervals to analyze the non-dominated frontier, extraction of complementary information, data labeling, and confusion matrix. To evaluate a-posteriori analysis, we calculate the accuracy with Eq. (5), presented in Sect. 3. We chose round 7 from range 3, with the lowest accuracy of 0.8102. To represent the best accuracy, we chose round 7 from range 5

**Table 2** The average, standard deviation, minimum and maximum accuracy values from 10 runs

| Range | Accuracy (ACC) |                    |         |         |
|-------|----------------|--------------------|---------|---------|
|       | Average        | Standard deviation | Minimum | Maximum |
| 1     | 0.9902         | 0.0017             | 0.9883  | 0.9925  |
| 2     | 1              | 0                  | 1       | 1       |
| 3     | 0.9082         | 0.0429             | 0.8102  | 0.9484  |
| 4     | 0.9960         | 0.0012             | 0.9942  | 0.9983  |
| 5     | 0.9760         | 0.0163             | 0.9575  | 1       |
| 6     | 1              | 0                  | 1       | 1       |



**Fig. 3** Non-dominated frontier of round 7 for the irradiance range 3 and 5

with a precision of 1. Although ranges 2 and 6 achieved the highest accuracy across all 10 runs, we did not choose them to demonstrate the following results. These two ranges obtained the same maximum probability value ( $p = 1$ ) for all conditions in the three portions, so we do not observe different probability information.

In Fig. 3, we present the non-dominated frontiers of round 7 for ranges 3 and 5. We can see that the objective functions are conflicting, and the irradiance ranges can reach different values in Compactness and Connectivity. We analyzed the correlations between Compactness and Connectivity and observed a negative correlation for almost all ranges. Only for interval 4 we observed a low correlation. Then, the objective functions are not conflicting, and using a mono-objective algorithm for this range would be enough.

Table 3 presents the additional information obtained for interval 3. We selected only some samples to demonstrate the information obtained in a-posteriori analysis. For each portion of the non-dominated frontier, we obtained different probability information based on the distribution of solutions. The columns No, Sc, and Oc of each portion correspond to the probability information of an unlabeled sample being assigned to the same grouping of the labeled samples of one of these conditions. The actual condition is the correct label for the samples. To analyze the performance of

**Table 3** Complementary information extracted from the three portions of non-dominated frontiers about the irradiance range 3 (893.252–911.252 W/m<sup>2</sup>) for the samples A to F

| Sample | Irradiance 893.252–911.252 |       |       |           |       |       |           |       |       | Actual Condition |
|--------|----------------------------|-------|-------|-----------|-------|-------|-----------|-------|-------|------------------|
|        | Portion 1                  |       |       | Portion 2 |       |       | Portion 3 |       |       |                  |
|        | No                         | Sc    | Oc    | No        | Sc    | Oc    | No        | Sc    | Oc    |                  |
| A      | 1                          | 0     | 0     | 0.923     | 0.077 | 0     | 0.598     | 0.402 | 0     | No               |
| B      | 0.533                      | 0.460 | 0.007 | 0.355     | 0.635 | 0.009 | 0.4933    | 0.505 | 0.002 | No               |
| C      | 0.062                      | 0.938 | 0     | 0.320     | 0.676 | 0     | 0.290     | 0.708 | 0.002 | Sc               |
| D      | 0.740                      | 0.260 | 0     | 0.364     | 0.635 | 0.001 | 0.290     | 0.708 | 0.002 | Sc               |
| E      | 0.223                      | 0.766 | 0.011 | 0.355     | 0.635 | 0.009 | 0.355     | 0.635 | 0.010 | Oc               |
| F      | 0.671                      | 0.329 | 0     | 0.718     | 0.282 | 0     | 0.573     | 0.423 | 0.004 | Oc               |

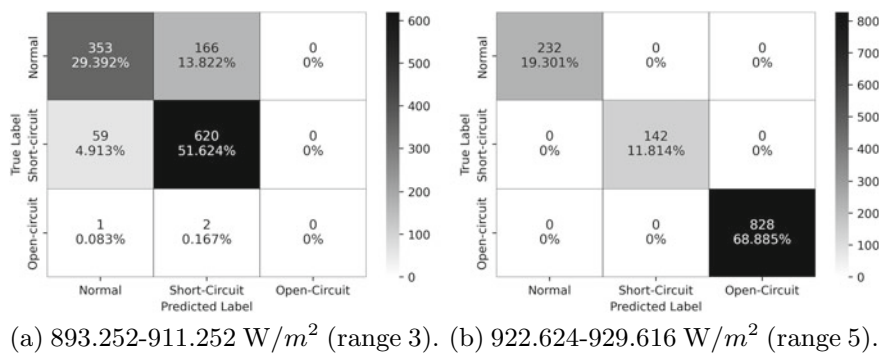
the algorithm in detecting failures, we considered as the final condition of the sample the maximum probability value obtained in one of the three portions of the frontier. In Table 3, for sample A, portion 1 indicates that the sample has a 100% chance of belonging to the same cluster of samples under normal conditions. For samples A and C, the condition suggested by a-posteriori analysis corresponds to the actual condition, but not for samples B, D, E, and F. The approach could not detect the faults for Open-circuit samples in this range. This difficulty may be associated with the small number of unlabeled and labeled samples of this condition in range 3.

Table 4 presents complementary information obtained for range 5. In this range, all unlabeled samples were correctly detected. Furthermore, for Open-circuit samples, the three portions reached the maximum probability ( $p = 1$ ), as we can see in samples E and F of Table 4. We observed the same results for ranges 1 and 4 regarding Open-circuit faults. All portions of the non-dominated frontier obtained homogeneous results and efficiently detected Open-circuit faults.

To analyze the overall performance of data labeling, we used the Confusion Matrix to understand the potential of the proposed approach in fault detection. In the matrix, “True Label” corresponds to the actual condition of the record and “Predicted Label” to the label extracted from complementary information of a-posteriori analysis. Figure 4a presents the Confusion Matrix for irradiance range 3, with an accuracy of 0.8102. For range 3, 166 normal data (13.822%), 59 Short-circuit faults (4.913%), and 3 Open-circuit faults were mispredicted. As mentioned earlier, we had a few unlabeled and labeled Open-circuit samples for this range. Figure 4b presents the Confusion Matrix of range 5, with an accuracy of 1. In this case, all the unlabeled samples were predicted correctly.

**Table 4** Complementary information extracted from the three portions of non-dominated frontiers about the irradiance range 5 (922.624–929.616 W/m<sup>2</sup>) for the samples A to G

| Sample | Irradiance 922.624–929.616 |              |          |           |       |          |              |              |          | Actual condition |
|--------|----------------------------|--------------|----------|-----------|-------|----------|--------------|--------------|----------|------------------|
|        | Portion 1                  |              |          | Portion 2 |       |          | Portion 3    |              |          |                  |
|        | No                         | Sc           | Oc       | No        | Sc    | Oc       | No           | Sc           | Oc       |                  |
| A      | 0.535                      | 0.465        | 0        | 0.560     | 0.440 | 0        | <b>0.858</b> | 0.142        | 0        | No               |
| B      | <b>1</b>                   | 0            | 0        | 0.979     | 0.021 | 0        | 0.858        | 0.142        | 0        | No               |
| C      | 0.312                      | <b>0.688</b> | 0        | 0.441     | 0.559 | 0        | 0.314        | 0.686        | 0        | Sc               |
| D      | 0.553                      | 0.447        | 0        | 0.560     | 0.440 | 0        | 0.334        | <b>0.666</b> | 0        | Sc               |
| E      | 0                          | 0            | <b>1</b> | 0         | 0     | <b>1</b> | 0            | 0            | <b>1</b> | Oc               |
| F      | 0                          | 0            | <b>1</b> | 0         | 0     | <b>1</b> | 0            | 0            | <b>1</b> | Oc               |



**Fig. 4** Confusion Matrix for the irradiance range 3 and 5. The values in the matrix represent the number of samples followed by the corresponding percentage of the total unlabeled data

5 Conclusion

This paper proposed MECA to capture the different clustering perspectives through objective functions. We proposed a-posteriori analysis to label data from the information of the non-dominated frontier. The frontier, divided into three portions, can evaluate which perspectives a cluster may have, providing information when we do not know the data distribution. Besides, a-posteriori analysis makes it possible to infer whether a cluster has more than one perspective (by analyzing the probabilities of the portions). The approach was applied in six irradiance ranges, and the results were validated through the accuracy and confusion matrix. A-posteriori analysis showed promising results regarding fault detection. In general, for all the irradiance ranges, the approach could detect the conditions of faults and normal data correctly, demonstrating a helpful approach to detecting and classifying faults in PV systems. The results also highlight the benefits of using non-dominated solutions to explore complementary information for data clustering. Relying on the irradiance range and

the type of fault, we observed that each portion of the frontier could better represent the data distribution, or we can have a homogeneous result between the three portions. This result indicates that using the division of the frontier instead of the entire non-dominated frontier could help better understand the data distribution. In future work, we would like to extend the proposed method to other problems requiring data labeling, such as semi-supervised learning problems.

**Acknowledgements** This work is financed by the ERDF—European Regional Development Fund, through the Operational Programme for Competitiveness and Internationalisation—COMPETE 2020 Programme under the Portugal 2020 Partnership Agreement, within project SmartPV, with reference POCI-01-0247-FEDER-068919.

## References

1. Ali, H.M.: Recent advancements in PV cooling and efficiency enhancement integrating phase change materials based systems-a comprehensive review. *Sol. Energy* **197**, 163–198 (2020)
2. Dhimish, M., Holmes, V., Mehrdadi, B., Dales, M., Mather, P.: Photovoltaic fault detection algorithm based on theoretical curves modelling and fuzzy classification system. *Energy* **140**, 276–290 (2017)
3. Zhu, H., Lu, L., Yao, J., Dai, S., Hu, Y.: Fault diagnosis approach for photovoltaic arrays based on unsupervised sample clustering and probabilistic neural network model. *Sol. Energy* **176**, 395–405 (2018)
4. Tina, G.M., Cosentino, F., Ventura, C.: *Monitoring and Diagnostics of Photovoltaic Power Plants*. Springer International Publishing (2016)
5. Mellit, A., Tina, G., Kalogirou, S.: Fault detection and diagnosis methods for photovoltaic systems: a review. *Renew. Sustain. Energy Rev.* **91**(February), 1–17 (2018)
6. Lazzaretti, A.E., da Costa, C.H., Rodrigues, M.P., Yamada, G.D., Lexinoski, G., Moritz, G.L., Oroski, E., de Goes, R.E., Linhares, R.R., Stadzisz, P.C., Omori, J.S., dos Santos, R.B.: A monitoring system for online fault detection and classification in photovoltaic plants. *Sens. (Basel, Switz.)* **20**(17), 4688 (2020)
7. Et-taleby, A., Boussetta, M., Benslimane, M.: Faults detection for photovoltaic field based on k-means, elbow, and average silhouette techniques through the segmentation of a thermal image. *Int. J. Photoenergy* **2020**, 1–7 (2020)
8. Cai, Y., Lin, P., Lin, Y., Zheng, Q., Cheng, S., Chen, Z., Wu, L.: Online photovoltaic fault detection method based on data stream clustering. *IOP Conf. Ser. Earth Environ. Sci.* **431**(1), 012,060 (2020)
9. Zhao, Y., Ball, R., Mosesian, J., de Palma, J.F., Lehman, B.: Graph-based semi-supervised learning for fault detection and classification in solar photovoltaic arrays. *IEEE Trans. Power Electron.* **30**(5), 2848–2858 (2015)
10. Zhu, X.J.: Semi-supervised learning literature survey. Tech. Rep. TR-1530, Dept. Comput. Sci., Univ. Wisconsin-Madison, Madison, WI, USA (2005)
11. Van Engelen, J.E., Hoos, H.H.: A survey on semi-supervised learning. *Mach. Learn.* **109**(2), 373–440 (2020)
12. Mahela, O.P., Shaik, A.G.: Recognition of power quality disturbances using s-transform based ruled decision tree and fuzzy c-means clustering classifiers. *Appl. Soft Comput.* **59**, 243–257 (2017)
13. Handl, J., Knowles, J.: An evolutionary approach to multiobjective clustering. *IEEE Trans. Evol. Comput.* **11**(1), 56–76 (2007)

14. Deb, K., Pratap, A., Agarwal, S., Meyarivan, T.: A fast and elitist multiobjective genetic algorithm: NSGA-II. *IEEE Trans. Evol. Comput.* **6**(2), 182–197 (2002)
15. MacQueen, J., et al.: Some methods for classification and analysis of multivariate observations. In: *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, vol. 1, pp. 281–297. Oakland, CA, USA (1967)
16. Kaufman, L., Rousseeuw, P.J.: *Finding Groups in Data: an Introduction to Cluster Analysis*. Wiley Series in Probability and Mathematical Statistics. Applied Probability and Statistics, vol. 344, pp. 68–125 (1990)
17. Kruskal, J.B.: On the shortest spanning subtree of a graph and the traveling salesman problem. *Proc. Am. Math. Soc.* **7**(1), 48–50 (1956)
18. Rampazzo, P.C.B., Yamakami, A., de França, F.O.: Evolutionary approaches for the multi-objective reservoir operation problem. *J. Control. Autom. Electr. Syst.* **26**(3), 297–306 (2015)
19. Mukhopadhyay, A., Maulik, U., Bandyopadhyay, S.: A survey of multiobjective evolutionary clustering. *ACM Comput. Surv. (CSUR)* **47**(4), 1–46 (2015)