



Deposited via The University of Leeds.

White Rose Research Online URL for this paper:

<https://eprints.whiterose.ac.uk/id/eprint/242034/>

Version: Accepted Version

Article:

Shiri, D., Shahmanzari, M. and Tanrisever, F. (2026) The online election campaign planning problem: Optimizing election campaign strategies with inaccurate information. Production and Operations Management. ISSN: 1059-1478

<https://doi.org/10.1177/10591478261438365>

This is an author produced version of an article published in Production and Operations Management, made available via the University of Leeds Research Outputs Policy under the terms of the Creative Commons Attribution License (CC-BY), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

Reuse

This article is distributed under the terms of the Creative Commons Attribution (CC BY) licence. This licence allows you to distribute, remix, tweak, and build upon the work, even commercially, as long as you credit the authors for the original work. More information and the full terms of the licence here:

<https://creativecommons.org/licenses/>

Takedown

If you consider content in White Rose Research Online to be in breach of UK law, please notify us by emailing eprints@whiterose.ac.uk including the URL of the record and the reason for the withdrawal request.

Davood Shiri

<https://orcid.org/0000-0003-2884-0047>

Leeds University Business School, Woodhouse, Leeds, LS6 1AN, UK

+44(0)113 343 8127

D.Shiri@leeds.ac.uk

Masoud Shahmanzari

<https://orcid.org/0000-0003-2019-4490>

Brunel Business School, Brunel University London, Kingston Ln, Uxbridge, London, UB8 3PH, UK

+44 (0)1895 265278

masoud.shahmanzari@brunel.ac.uk

Fehmi Tanrisever*

<https://orcid.org/0000-0002-3921-3877>

Faculty of Business Administration, Bilkent University, Ankara, Turkey

+90 312 2901276

tanrisever@bilkent.edu.tr

*Corresponding author: tanrisever@bilkent.edu.tr

The Online Election Campaign Planning Problem: Optimizing Election Campaign Strategies with Inaccurate Information

Abstract: Effective management of election campaigns involves dynamic decision-making under uncertainty. Traditional approaches rely heavily on pre-planned strategies that often fail to adapt to real-time changes in voter sentiment and external factors. This paper introduces the Online Election Campaign Planning Problem (OECPP) to optimize the scheduling of campaign activities in the context of U.S. presidential elections. OECPP incorporates sequentially updated predictions that represent assessments of the impact of campaign activities over the course of the campaign. Since these predictions evolve in response to new information and their accuracy cannot be fully assessed without perfect information, we develop deterministic and randomized online algorithms for OECPP that can operate effectively under unreliable and evolving predictions.

We evaluate the performance of our algorithms using the competitive ratio (CR), a metric particularly useful when probabilistic modeling is impractical. We begin by establishing a tight upper bound on the CR of the online algorithms for the OECPP under unreliable reward predictions. We then introduce a sequential setup-based CR metric to capture the value of reoptimization as new predictions arrive, and we design deterministic and randomized algorithms that are optimal under this metric. Using data from U.S. presidential elections, we show that randomized online algorithms can significantly outperform their deterministic counterparts in terms of empirical CR. We also find that the effectiveness of randomized algorithms is driven by two factors: the selection of prediction samples for generating activity scenarios and the randomization cut-off, which determines the scenarios to be randomized. The benefit of randomization is non-monotonic, and the best empirical CR is achieved by selectively adding prediction samples to the randomization set.

Key words: Online Optimization, Competitive Ratio, Election Logistics, Inaccurate Information, Data Analytics

Date received 16 September 2024; accepted 07 March 2026 after two revisions

Handling Editor: Stefanus Jasin

1 Introduction

Technological advancements in data collection and analysis have revolutionized campaign planning across diverse fields such as electoral politics, advertising, and disaster relief. Over the past few decades, there has been a marked transition from the traditional methods of meticulous pre-planning and rigid campaign execution to a more dynamic, flexible approach driven by real-time data and predictive analytics (Rosário and Dias 2023). This evolution has provided decision makers and campaign managers with unprecedented foresight into operational uncertainties, significantly enhancing the decision-making process during the critical phases of campaign planning and execution.

In tandem with practical advancements, academic interest in online adaptations of traditional planning problems has grown substantially. These models provide a common framework for *selective* decision-making problems such as trip planning (Tlili and Krichen 2021), election campaign planning (Shahmanzari et al. 2020), and nurse routing (Cinar et al. 2021). A common assumption in these problems is that the rewards for visiting specific locations are deterministic and known in advance. However, this assumption only holds in certain real-life scenarios where preferences and conditions are static and do not change dynamically. For example, in the context of election campaigns, the strategic value of visiting a city can shift in response to actions by the opposition or new intelligence, making pre-campaign reward estimations highly unreliable, posing a challenge for planners to adapt without relying on predefined probability distributions.¹

Motivated by the planning and execution of U.S. presidential election campaigns, this manuscript introduces the Online Election Campaign Planning Problem (OECPP).² The complexity of planning and executing election campaigns encompasses numerous decisions and factors under uncertainty. Political parties initiate preparations months ahead, with candidates visiting multiple states within a constrained framework of time and geography. The process involves intricate planning of the speech calendar and events to optimize exposure and impact.

The challenge of election campaigning is further compounded by the uncertainty surrounding the impact of campaign activities on voter decisions—referred to as the *reward* in this context. In our study, this reward specifically refers to the *contribution* of particular campaign activities, conducted at specific locations and times, to influencing voter choices. Reward predictions are notoriously difficult, as they are affected by fluctuating polls, sampling biases, methodological shortcomings, and unforeseen political events that occur in real time (Ansolabehere and Hersh 2012). A clear example of this unpredictability was seen during the 2016 U.S. presidential election campaign. Despite leading in the polls shortly before Election Day, Hillary Clinton’s decision to focus campaign efforts on Arizona instead of on battlegrounds such as Wisconsin illustrates how campaign strategies,³ shaped by uncertain reward expectations, can lead to unexpected election outcomes (Devine 2018).

¹ An example of a campaign adapting its schedule after an unforeseen event occurred during the recent 2024 U.S. presidential election when Donald Trump faced an attempted assassination at a rally in Pennsylvania. Following the incident, his campaign rapidly adjusted its plan. The candidate shifted his focus to Pennsylvania and neighboring states to leverage the event to rally his base and present an image of resilience and strength (<https://www.cbsnews.com/pittsburgh/news/trump-rally-butler-pennsylvania/>). The dramatic moment altered his travel plans to maximize political gains from the incident. Similarly, during the same election, a devastating Category 5 hurricane struck key regions in the U.S. just weeks before election day. Both campaigns updated their schedules to address the disaster. Kamala Harris focused on visiting affected areas to emphasize her administration’s climate change initiatives (<https://www.nytimes.com/2024/10/05/us/politics/kamala-harris-north-carolina-hurricane-helene.html>).

² We note that our model and insights are generally applicable to other elections (e.g., primary elections) and campaign planning scenarios (e.g., humanitarian logistics).

³ <https://www.theguardian.com/us-news/2016/nov/02/hillary-visits-arizona-in-unprecedented-show-of-bravado-for-democratic-nominee>

The unreliability of reward predictions has been further highlighted by inaccuracies in polling data observed in practice, as these discrepancies often reflect the inherent challenges in predicting voter behavior and campaign effectiveness. In the U.K.'s 1992 general election, polling inaccuracies led to a significant misjudgment of voter behavior. Pre-election polls suggested a tight race between the Conservative Party and the Labour Party. However, when the election results were announced, the Conservatives secured an unexpected victory contrary to the polling predictions (Jowell et al. 1993). Similarly, in the U.S. 2020 presidential election, polling projections faced scrutiny for their inability to accurately predict the final result. Many polls indicated a significant lead for the Democratic candidate, Joe Biden, over the Republican candidate, Donald Trump. However, the election results revealed that Trump outperformed expectations, with Biden's actual margin of victory being considerably narrower than polls had suggested. The polling errors raised questions about the reliability of surveys and the challenges in capturing the evolving political dynamics. These examples highlight the inherent uncertainties in reward predictions, as unforeseen factors can often influence electoral outcomes in ways that defy conventional expectations (Campbell 2024).

To address these limitations, we model campaign planning under *unknown* rewards. Specifically, we assume that the reward for an activity at a given location–time pair is unknown *ex ante* and would be revealed only under perfect information—which is typically unavailable in elections—so the true rewards may never be disclosed to the decision maker. Instead, the decision maker observes a prediction matrix whose entries are dynamically updated and may be unreliable. We therefore formulate an online variant of the election campaign planning problem and develop online methods tailored to this evolving information. Accordingly, we ask: (1) how should the campaign plan be reoptimized as unreliable predictions arrive dynamically, and (2) how can randomization—i.e., making controlled probabilistic choices rather than fixed decisions in the campaign plan—improve robustness to prediction errors?

For online optimization problems where stochastic or robust optimization techniques are not applicable, due to the impracticality of probabilistic evaluation or assessment of uncertainty sets, existing literature recommends employing the competitive ratio (CR) as a metric to gauge decision quality (Jaillet and Wagner 2008, Manshadi et al. 2012, Golrezaei et al. 2014, Ma and Simchi-Levi 2020, Aouad and Saban 2023). Accordingly, we will adopt this metric to assess the performance of our algorithms and models. The contribution of our work can be summarized as follows:

- First, inspired by the U.S. presidential election campaigns, we introduce the OECPP with unknown rewards that are not associated with a prior probability distribution.
- Second, we adopt the setup-based CR criterion from Ma and Simchi-Levi (2020) and prove a tight upper bound on the setup-based CR of online algorithms, along with an online algorithm that has a matching setup-based CR.

- Third, we extend the setup-based CR criterion to a sequential setting to incorporate the effect of sequentially received information (i.e., predictions) into the performance metric. In this way, we compare the effect of reoptimization in response to newly observed information and quantify the impact of different reoptimization policies. Specifically, we introduce deterministic and randomized online reoptimization policies that are proven to be the best possible under the sequential setup-based CR metric.
- Fourth, to validate our results, we investigate a real-life application of the OECPP to election logistics as a case study, focusing on the U.S. presidential elections. We find that using properly randomized online algorithms can significantly improve the empirical CR of online algorithms. We show that randomized online algorithms in our case can outperform the best deterministic online algorithms.
- Fifth, we also identify two key parameters that drive the effectiveness of randomized algorithms: (1) the set of prediction samples to be inputted into the reoptimization formulation to generate different activity scenarios for the politician, and (2) randomization cut-off to determine which scenarios must be selected for randomization. We find that the benefit of randomization is non-monotonic in the randomization cut-off, meaning that it is neither optimal to set it too high nor too low. Additionally, selectively adding prediction samples to the randomization set results in the best empirical CR.

The remainder of the paper is organized as follows. We review the related literature in Section 2, and formally describe the problem in Section 3. Next, in Section 4, we present the preliminaries for online analysis and develop online reoptimization and randomization strategies under dynamically updated, unreliable predictions. We investigate a case study of the OECPP in election logistics along with computational results in Section 5. We provide managerial insights and a conclusion in Section 6, and discuss limitations and future research directions in Section 7. All proofs of our theoretical results are presented in the E-Companion in Section EC.3.

2 Literature review

The OECPP is a variant of the classical Roaming Salesman Problem (RSP), a problem in which the goal is to maximize the total collected reward by visiting locations under time constraints while balancing routing and scheduling decisions. The OECPP focuses specifically on the scheduling aspects of decision-making and eliminates the routing components of the RSP. Furthermore, it incorporates election-related constraints, such as mandated rest days. Unlike the RSP, the OECPP does not assume prior access to an accurate time-dependent reward function. Instead, it introduces a sequential mechanism in which reward predictions with unknown accuracy are revealed in real time. The RSP has been well-studied in the operations research literature in the absence of uncertainty (Shahmanzari et al. 2020, Shahmanzari and Aksen 2021, Shahmanzari and Mansini 2024); however, to the best of our knowledge, there is no research addressing this problem in the presence of uncertainty. Since our study focuses on modeling the problem under real-time uncertainty

in the context of political elections, our literature is organized around two strands. First, we cover relevant online optimization research; second, we review the political science literature that informs our modeling choices.

2.1 Online optimization

Over the past few decades, various online optimization problems have been studied. In much of this literature, *actual* information about uncertain parameters is sequentially revealed to the decision maker over time (Ausiello et al. 2001, Jaillet and Wagner 2008, Golrezaei et al. 2014, Ma and Simchi-Levi 2020, Aouad and Saban 2023). This literature uses CR as the performance criterion to analyze these online optimization problems. The consensus in designing the best algorithms from a CR perspective involves utilizing exact mathematical formulations that incorporate partially available information to update decisions sequentially. It is important to note that these studies assume that actual information about uncertain parameters is gradually revealed to the decision maker, allowing for more informed decisions due to the guaranteed accuracy of each piece of sequentially revealed information. In contrast, in the OECPP, the actual rewards remain unknown, and only predictions are provided to the decision maker.

In a recent stream of research, there has been growing interest in studying online optimization problems with unreliable predictions (Lykouris and Vassilvitskii 2021, Gouleakis et al. 2023, Rutten and Mukherjee 2024, Shiri et al. 2024), where online algorithms are provided with predictions in a black-box manner, accompanied by unknown errors. The primary objective of this line of research is to develop online algorithms that perform close to optimal when the predictions are accurate, while also maintaining theoretical guarantees in terms of worst-case CR, even in scenarios with highly erroneous predictions. A key goal is to create an online algorithm with three key properties: (1) consistency, ensuring close alignment with the offline optimal solution under flawless predictions; (2) smoothness, ensuring controlled degradation of solution quality as prediction errors increase, while still maintaining performance relative to prediction accuracy and the offline optimal solution; and (3) robustness, which limits the algorithm's output within a predetermined range relative to the offline optimal solution regardless of prediction accuracy. Most work in this area has focused on theoretical investigations from a worst-case CR viewpoint, where the CR is modeled as a function of prediction error (Lykouris and Vassilvitskii 2021, Rutten and Mukherjee 2024, Shiri et al. 2024). Similar to this line of work, we incorporate unreliable predictions into our competitive analysis of the OECPP, providing our CRs as functions of prediction error to measure the consistency, smoothness, and robustness of our algorithms simultaneously.

In the studies above, predictions are revealed to the decision-maker only once at the beginning, while the actual information is revealed sequentially. In contrast, in the OECPP, predictions are updated sequentially (i.e., revised over time) and the actual rewards remain unknown. This key difference distinguishes our analysis from the existing literature on online optimization with predictions.

2.2 Political science on campaign planning

Our study aims to determine the optimal number and sequence of activities during an election campaign to maximize the resulting reward. To our knowledge, the political science literature has not explicitly addressed how parties optimize the timing and sequencing of campaign activities. However, it provides evidence on three key factors that influence the *reward* of such activities: (1) state-specific rewards assigned to conducting campaign activities in different states, reflecting their relative importance to the decision-maker; (2) timing effects over the campaign calendar; and (3) diminishing returns to increasing the number of activities in a state. As detailed below, these findings inform our modeling choices.

Numerous studies in the political science literature have examined the state-specific rewards associated with campaign activities. A primary determinant of location importance is electoral competitiveness. Swing states, where neither political party has a significant advantage, receive disproportionate attention from candidates (Beachler et al. 2015, Hill and Huber 2017). For example, in recent U.S. presidential elections, both major candidates have concentrated a large share of their visits in a small set of swing states. High-population areas with more electoral votes are also prioritized due to their larger electorates (Shaw 2008, Belenky 2016).

Taken together, this evidence is consistent with a simple and widely discussed targeting heuristic in the literature: allocate candidate appearances to states where the *marginal electoral payoff* is plausibly highest, meaning where an additional day of effort is most likely to change the electoral outcome. In an Electoral College setting, this marginal payoff increases with the electoral prize at stake (electoral votes) and with electoral competitiveness (close races), so formal models imply disproportionate attention to large, highly competitive states (Brams and Davis 1974, Colantoni et al. 1975). Empirical analyses that use candidate appearances as a proxy for campaign effort report patterns consistent with concentrating resources in electorally consequential states (Bartels 1985). Accordingly, this heuristic is often summarized as prioritizing states that are both electorally valuable (high electoral vote stakes) and electorally competitive (close races), while deprioritizing predictable partisan strongholds. This rule-of-thumb can be viewed as a descriptive baseline for *where* campaigns tend to concentrate effort; our contribution is to move beyond such static triage logic by formalizing the *day-by-day* scheduling problem under (i) timing effects over the campaign calendar, (ii) diminishing returns to repeated activities within a state, and (iii) feasibility constraints such as mandated rest and visit-spacing, while (iv) explicitly modeling that campaigns rely on sequentially updated *unreliable* reward predictions.

Turning to the timing of visits, evidence from political science literature suggests that campaign activities tend to have a greater impact during the *early* and *late* stages of the campaign. Holbrook (1996) argues that events occurring early in the campaign have a significant impact on the candidate's public image and help set the agenda for the remainder of the campaign. Similarly, Finkel (1993) and Broockman and Kalla (2023)

discuss how early campaign appearances play a significant role in solidifying voters' initial preferences and securing weak partisan identifiers. On the other hand, Herr (2002) show that late-stage appearances by Bill Clinton after October 1, 1996, were positively associated with increasing votes for him and reducing support for his opponents, highlighting the importance of late-stage campaign activities. In line with this, Brox and Giammo (2009) argue that late-stage campaign engagements can have a significant impact, particularly in swaying undecided voters. In the final days of the 2024 U.S. presidential election—one of the tightest races in modern history, with national polls showing a 47%-47% tie between Kamala Harris and Donald Trump⁴—the importance of late-stage campaign activities became especially pronounced, as even marginal shifts among undecided voters in battleground states like Pennsylvania and Georgia had the potential to tip the electoral outcome.⁵

Regarding the number of campaign activities, Shaw (1999) and Herr (2002) confirm that the number of campaign appearances is positively correlated with vote shares for candidates. However, it is notable that the positive effect of multiple engagements on voters does not scale proportionally with the number of appearances due to voter fatigue (Broockman and Kalla 2023).⁶ That is, increasing the number of campaign activities within a state yields diminishing marginal returns.

3 Problem

In the U.S., the winner of the presidential election is determined through an indirect election process known as the Electoral College system. In total, there are 538 electors in the Electoral College, and the candidate who wins a majority of electoral votes—at least 270—is elected president. Although citizens vote for their preferred presidential candidate on Election Day, they are actually casting votes for their state's chosen slate of electors. Afterward, the electors meet in their respective states to formally cast their votes for president, and the candidate with a majority of electoral votes becomes president. In this context, presidential campaigns play a critical role in shaping voter preferences and turnout, particularly in competitive states where outcomes are unpredictable. Typically, candidates campaign by visiting key states over a period of about two months, beginning shortly after their party conventions and continuing until Election Day. In this section, we first elaborate on the online aspect of election campaign planning, then present the formal notation and problem description. Finally, we provide an exact mathematical formulation that solves the OECP problem optimally, assuming access to perfect information.

⁴ <https://edition.cnn.com/2024/10/25/politics/cnn-poll-harris-trump>

⁵ <https://www.theguardian.com/us-news/2024/oct/13/harris-trump-election-race-campaign>

⁶ <https://apnews.com/article/donald-trump-crowd-size-beyonce-kamala-harris-dfba0b3cdfa3b92009d7ed7bc938157e>

3.1 The Online Election Campaign Planning Problem

The OECPP involves planning an election campaign by visiting a subset of states⁷ $I = \{1, 2, \dots, n\}$ to collect rewards over a planning horizon represented by the set $T = \{1, 2, \dots, \tau\}$, where τ denotes the number of days in the horizon. The realized reward from conducting an activity in state $i \in I$ on day $t \in T$ depends on the cumulative number of activities conducted in state i over the entire campaign (Shaw 1999, Herr 2002, Broockman and Kalla 2023). Accordingly, let $S = \{0, 1, 2, \dots, \tau\}$ denote the potential total activity counts for a state over the campaign. We refer to each $s \in S$ as an activity-count *scenario* and index rewards by this count.

The objective of the OECPP is to maximize collected rewards from states across all days of the campaign. These rewards are conceptualized as the *contribution* of conducting an activity in each state $i \in I$ on each day $t \in T$, given the scenario $s \in S$, to increase the party's share of the electoral vote; i.e., we use the term *contribution* to specifically focus on the reward that comes directly from an activity in a certain state on a certain day for a given scenario. Notably, a state with a large number of voters may not necessarily be considered a high-reward state in our model if its voting outcome is already predictable. For example, in partisan strongholds such as California (Democratic) or Alabama (Republican), campaign efforts have limited marginal impact on electoral outcomes.

Building on the political science literature presented in Section 2.2, and capturing the first-order issues in the context of the U.S. presidential election campaign, our model associates each state $i \in I$ with an unknown reward characterized by: (1) a base reward function (θ_i), determined by state-specific properties; (2) a time-dependent factor (Φ_t), which influences the reward based on the timing of the activity; and (3) a scenario-dependent factor (Ξ_s), which influences the reward depending on the number of activities conducted in state i over the entire campaign. Guided by the political science literature and real-world election dynamics, Section 5.2 and E-Companion EC.4 specify the functional forms of these components and the calibrated parameter values used in our empirical benchmarking.

Accordingly, we denote the actual reward from conducting an activity in state i on day t in scenario s by $\theta_{its} = \theta_i \times \Phi_t \times \Xi_s$ for $i \in I, t \in T, s \in S$. Hereafter, we let $\Theta = \{\theta_{its} : i \in I, t \in T, s \in S\}$ be the global offline reward matrix. Notably, θ_{its} is unknown to the politician due to inherent uncertainty. Moreover, as the campaign unfolds, the decision maker's perception of these rewards may shift in response to new information. For example, a state that was initially insignificant for the party may become highly important if the opposing party delivers a critical speech there. The absence of historical data on *contributions* from campaign activities in different states to the victory of different political parties in the U.S. presidential

⁷ The term "state" refers to campaign locations within the context of the U.S. presidential election. In other electoral settings, this term may correspond to cities or other geographic units. In our study, each state is treated as a single decision unit. While we acknowledge that Maine and Nebraska allocate electoral votes by congressional district, they are aggregated here for modeling simplicity. We justify this abstraction by noting that, in practice and in all our experiments, these states are not visited; hence, the distinction does not affect the results or conclusions.

elections makes it impossible to establish a suitable probability distribution or uncertainty set. Hence, the absence of precise information about the time-dependent rewards of each state, along with their real-time fluctuations during the campaign, turns election campaign planning into an online optimization problem.

As noted above, θ_{is} is not known to the politician in the OECPP. Instead, the politician relies on a set of predictions (beliefs) that are dynamically updated each day by a predictive model, such as a machine learning predictor that draws on multiple sources, including expert opinions and media, to generate time-dependent reward predictions for upcoming activities. That is, the predictive engine updates its belief about θ_{is} based on real-time information at the beginning of the day λ , where λ belongs to the set $\Lambda = \{-k + 1, \dots, \tau - k\}$ and k denotes the length of the scheduling freeze window, formally defined in Remark 2. We denote the prediction released on day λ for the reward of performing an activity in state i on day $t \in T$ in scenario $s \in S$ by $\hat{\theta}_{is\lambda}$. It is important to note that each political party relies on its own predictive model, constructed from its available resources and information. However, the exact quantification of the error term in these predictions remains elusive, as the true rewards of campaign activities cannot be precisely measured—even after the election outcome is known. Hereafter, we let $\hat{\Theta} = \{\hat{\theta}_{is\lambda} : i \in I, t \in T, s \in S, \lambda \in \Lambda\}$ be the global prediction matrix that is revealed to the politician over the course of the election campaign. For a detailed example of $\hat{\Theta}$ and $\hat{\theta}_{is\lambda}$, we refer the reader to Section 4.6.1 and Table 2.

REMARK 1. We emphasize that our models and algorithms are directly compatible with any predictive model. This flexibility is a key advantage of our algorithms, as they are agnostic to the specific components of the predictive model.

3.2 The deterministic components of the OECPP

We next detail the deterministic components of the OECPP as follows. As discussed, $I = \{1, 2, \dots, n\}$ is the set of states and $T = \{1, 2, \dots, \tau\}$ is the set of days in the planning horizon. To reflect operational limits, we set the politician's maximum number of activities per day to α . Furthermore, inspired by real-life election campaigns—where candidates typically avoid scheduling daily activities in the early weeks but campaign daily in the final two weeks—our model enforces at least β rest days in any seven consecutive days of the campaign, except during the final two weeks. Lastly, based on historical U.S. election data, politicians typically avoid conducting campaign activities in the same state multiple times within a short time interval to prevent voter fatigue and message saturation. Accordingly, our model limits the number of activities in any given state to at most $\bar{\beta}$ within any seven consecutive days.

For U.S. presidential elections, party leaders often travel via private jets or chartered transportation, which allows them to move flexibly between locations, so we do not explicitly consider routing. Similarly, although incorporating a budget constraint may seem intuitive, its practical relevance is limited due to uncertainty in budget compliance,⁸ particularly because logistical expenses constitute only a small portion of overall

⁸ <https://www.newsnationnow.com/politics/2024-election/harris-campaign-millions-debt/>

costs in U.S. presidential election campaigns.⁹ This allows us to focus on first order issues that are central to campaign planning, such as the timing and number of visits.¹⁰

3.3 Offline formulation

We develop a mathematical model for the offline variant of OECPP, a multi-period Integer Linear Programming (ILP) formulation, which provides the optimal schedule for the politician with access to Θ . Specifically, the offline model is based on θ_{its} for $i \in I$, $t \in T$, and $s \in S$. The formulation for the offline version is presented in (1) - (10). In this formulation, we track activities using a binary decision variable X_{its} , which denotes whether an activity is conducted in state $i \in I$ on day $t \in T$, given that state i has been visited $s \in S$ times during the campaign.

REMARK 2. Recognizing that candidates often have limited flexibility to alter immediate activity plans due to logistical constraints, pre-arranged commitments, and media announcements, campaign schedules must be finalized several days in advance, a period we refer to as the scheduling freeze window.¹¹ To reflect this constraint, we assume that the optimization corresponding to each campaign day $t \in T$ is performed k days prior to day t , where k denotes the length of the scheduling freeze window (e.g., the offline formulation is solved k days before Day 1).

$$\max \sum_{s \in S} \sum_{i \in I} \sum_{t \in T} X_{its} \theta_{its} \quad (1)$$

s.t.

a) Coupling decision variables:

$$X_{its} \leq R_{is}, \quad i \in I, t \in T, s \in S \quad (2)$$

$$X_{its} \leq Z_{it}, \quad i \in I, t \in T, s \in S \quad (3)$$

$$X_{its} \geq R_{is} + Z_{it} - 1, \quad i \in I, t \in T, s \in S \quad (4)$$

$$\sum_{s \in S} R_{is} = 1, \quad i \in I \quad (5)$$

$$\sum_{t \in T} Z_{it} = \sum_{s \in S} s R_{is}, \quad i \in I \quad (6)$$

b) Modeling rest requirements:

$$\sum_{j=t}^{t+6} V_j \geq \beta, \quad 1 \leq t \leq \tau - 20 \quad (7)$$

$$\sum_{i \in I} Z_{it} \leq \alpha(1 - V_t), \quad t \in T \quad (8)$$

⁹ <https://www.voanews.com/a/usa-us-politics-billions-spent-elect-us-presidents/6192568.html>

¹⁰ While other practical constraints could be included, the model is kept generic to focus on first-order issues, given the limited availability of detailed information about each party's logistical needs.

¹¹ <https://hls.harvard.edu/bernard-koteen-office-of-public-interest-advising/a-quick-guide-to-working-on-political-campaigns/>

c) *Practical constraints tailored for U.S. elections:*

$$\sum_{j=t}^{t+6} Z_{ij} \leq \bar{\beta}, \quad i \in I, 1 \leq t \leq \tau - 6 \quad (9)$$

d) *Decision variable definition:*

$$R_{is}, X_{its}, Z_{it}, V_t \in \{0, 1\} \quad (10)$$

The objective function, provided in (1), maximizes the summation of collected rewards. Constraints (2)-(6) couple binary decision variables. Z_{it} indicates whether the reward associated with state $i \in I$ is collected by the politician on day $t \in T$, while R_{is} denotes whether state $i \in I$ hosts $s \in S$ activities throughout the campaign.¹² Constraints (7-8) capture the rest requirements by defining V_t as a rest-day indicator and enforcing at least β rest days in any seven-day window (excluding the last two weeks), while limiting the total number of activities on a given day to at most α . Constraints (9) ensure that each state can be visited at most $\bar{\beta}$ times within any seven-day window. Constraints (10) define the binary decision variables. Table 1 presents the key notation for the model.

Table 1 Notation

Notation	Description
Sets	
I	Set of states (decision units), i.e., $I = \{1, \dots, n\}$.
T	The set of τ days comprising the campaign planning horizon, i.e., $T = \{1, \dots, \tau\}$.
S	The set of possible values for the number of activities conducted in the same state during the campaign, i.e., $S = \{0, 1, 2, \dots, \tau\}$.
Θ	The global offline reward matrix, i.e., $\Theta = \{\theta_{its} : i \in I, t \in T, s \in S\}$.
$\hat{\Theta}$	The global prediction matrix, i.e., $\hat{\Theta} = \{\hat{\theta}_{its\lambda} : i \in I, t \in T, s \in S, \lambda \in \Lambda\}$.
Γ	The set of observable deterministic input parameters at the beginning of the campaign, i.e., $\Gamma = \{I, T, \alpha, \beta, \bar{\beta}, k\}$.
Parameters	
α	Maximum number of activities allowed in each day.
β	Minimum number of rest days per seven consecutive days (excluding the final two weeks).
$\bar{\beta}$	Maximum number of days within any seven consecutive days during which an activity is conducted in the same state.
k	Length parameter of the scheduling freeze window. After each optimization, the schedule for the next k campaign days remains fixed and cannot be revised.
$GE \geq 0$	Global prediction error (unknown).
Decision variables	
Z_{it}	Binary variable indicating if the politician holds an activity in state $i \in I$ on day $t \in T$ and collects the associated reward.
R_{is}	Binary variable indicating if state $i \in I$ hosts s activities throughout the campaign.
X_{its}	Binary variable indicating if state $i \in I$ accommodates an activity on day $t \in T$ while hosting s activities throughout the campaign.
V_t	Binary variable indicating if day $t \in T$ is a rest day.

¹² We note that $s = 0$ is included in S for notational completeness. Constraints (2)–(6) jointly imply that $X_{i0} = 0$ for all $i \in I$ and $t \in T$, so this case never contributes to the objective.

4 Theoretical analysis

With the problem formally defined in Section 3, we now turn to the theoretical analysis in Sections 4.1–4.6. We begin by introducing the setup-based CR as a baseline worst-case performance benchmark, and then define a prediction-error metric that allows us to formulate performance guarantees for online algorithms as a function of prediction error. Because campaigns receive new predictions throughout the horizon, we then extend this notion to a sequential setup-based CR that accounts for reoptimization. Under this sequential framework, we identify the optimal deterministic reoptimization policy and propose a randomized reoptimization policy that allocates choices probabilistically among multiple high-quality schedules.

4.1 Performance criterion

The conventional performance measure for analyzing online algorithms is the CR criterion (Jaillet and Wagner 2008, Golrezaei et al. 2014, Ma and Simchi-Levi 2020, Manshadi et al. 2022, Aouad and Saban 2023, Feng and Niazadeh 2025). Online algorithms are broadly categorized as either deterministic or randomized. Deterministic online algorithms consistently produce the same output when applied to the same problem instance multiple times. In contrast, randomized online algorithms can yield different solutions for the same problem instance in separate runs. Accordingly, we evaluate randomized algorithms using expected objective values.

Inspired by the definition of *setup* in Ma and Simchi-Levi (2020), we evaluate the performance of online algorithms based on the setup of deterministic input parameters realized prior to the start of the campaign, denoted as Γ in Table 1. This setup-based approach to CR is applied within the context of our problem.

DEFINITION 1. For any deterministic or randomized algorithm labeled by ALG and any setup Γ , the setup-based CR (SCR_Γ) is defined as

$$0 \leq \inf_{\hat{\Theta}, \Theta} \frac{\mathbb{E}[ALG(\Gamma, \hat{\Theta}, \Theta)]}{OPT(\Gamma, \hat{\Theta})} \leq 1,$$

where $\mathbb{E}[ALG(\Gamma, \hat{\Theta}, \Theta)]$ represents the (expected) *actual* objective function value of the online algorithm with respect to Γ (which is known prior to the start of the campaign), $\hat{\Theta}$ (which is input to ALG day-by-day and affects its decisions), and Θ (which remains unknown to ALG but is used to compute its actual total collected reward). Furthermore, $OPT(\Gamma, \hat{\Theta})$ is an upper bound on the objective function of the offline optimum presented in Lemma EC.1.

The definition above provides a guarantee on the expected objective function value of the online algorithm given the setup Γ , against an upper bound on the offline optimum which is obtained with full access to $\hat{\Theta}$ from the beginning, i.e., the upper bound on the offline optimum which is computed using $\hat{\Theta}$ holds for any Θ .

REMARK 3. We acknowledge that Definition 1 could be refined by replacing $OPT(\Gamma, \hat{\Theta})$ with $OFF(\Gamma, \Theta)$, the exact objective function value of the offline optimum calculated with full access to Θ from the beginning. However, this makes the setup-based CR analysis difficult, so we upper bound $OFF(\Gamma, \Theta)$ by $OPT(\Gamma, \hat{\Theta})$ in our definition. We note that in CR analysis, it is standard practice to upper bound the objective function value of the offline optimum, as seen in the literature (Golrezaei et al. 2014, Ma and Simchi-Levi 2020, Wang et al. 2022). We present our upper bound on $OFF(\Gamma, \Theta)$ in Lemma EC.1.

To render this worst-case benchmark informative with respect to prediction accuracy, we next introduce an explicit prediction error metric and subsequently formulate our CR analysis as a function of it.

4.2 Incorporating prediction accuracy into the CR

Analyzing the CR in relation to prediction accuracy is a novel approach in the online optimization literature. Different authors have employed various definitions of prediction error (i.e., prediction accuracy) based on the specific characteristics of the online optimization problems they study (Gouleakis et al. 2023, Rutten and Mukherjee 2024, Shiri et al. 2024). Inspired by the literature, to incorporate the prediction accuracy into our analysis, we define the following parameter which remains unknown to the online algorithm.

DEFINITION 2. The global normalized prediction error (GE) is defined as the maximal absolute percentage error for the prediction matrix $\hat{\Theta}$, i.e.,

$$GE = \max_{i \in I, t \in T, s \in S, \lambda \in \Lambda} \left| \frac{\theta_{its} - \hat{\theta}_{its\lambda}}{\theta_{its}} \right|.$$

The parameter GE is unknown to the online algorithm. This is due to the fact that the actual rewards cannot be observed without access to perfect information (see Section 3.1). We solely define GE to establish a meaningful quantitative measure for prediction accuracy. That is, we do not assume that the value of this parameter is known, but we respect its existence and incorporate it as an unknown parameter in our analysis. Next, we employ GE to analyze the setup-based CR of online algorithms for the OECPP.

4.3 An online algorithm with the best possible SCR_{Γ}

Using GE , we derive in Theorem 1 a tight upper bound on the SCR_{Γ} of online algorithms and, in Proposition 1, present a simple online algorithm that attains this bound.

THEOREM 1. For any setup Γ and any (deterministic or randomized) online algorithm ALG , there exists a $\hat{\Theta}$ and a Θ which enforce the SCR_{Γ} of $\max\{0, \frac{1-GE}{1+GE}\}$ on ALG for the OECPP.

Figure 1 visualizes the upper bound in Theorem 1 where GE ranges from 0 to 1. It can be observed that when $GE = 0$, the upper bound is 1, implying that an optimal online algorithm can achieve the best possible solution with access to perfect information in an extreme case where the global error is zero. As the error increases, the upper bound moves towards zero from one. When GE exceeds a certain threshold (i.e., $GE \geq 1$), the upper bound equals zero.

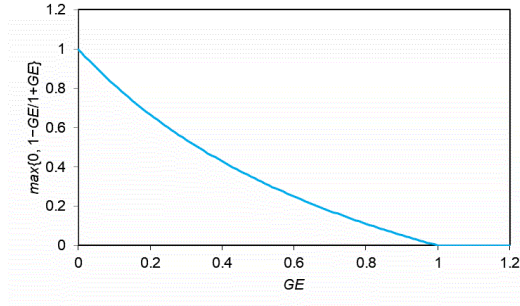


Figure 1 Upper bound on the setup-based CR of online algorithms for the OECPP across different values of GE . Because the upper bound is presented as a mathematical function of the error metric, the figure also shows consistency, smoothness, and robustness discussed in Gouleakis et al. (2023), Rutten and Mukherjee (2024), Shiri et al. (2024), see Section 2.1 for more details.

PROPOSITION 1. *For any Γ , a deterministic algorithm that adheres to the optimal solution of the exact ILP in Section 3.3 applied to the first revealed predictions attains an SCR_{Γ} which is at least equal to the upper bound in Theorem 1. We refer to this algorithm as the single-shot formulation-based (SFB) algorithm.*

This result demonstrates the tightness of the upper bound presented in Theorem 1. The key insight in this analysis is the necessity of using an exact mathematical formulation applied to the predictions with unknown errors to achieve the best policy in the worst-case for the politician. That is, a policy that does not use the exact mathematical formulation, and hence does not guarantee solving the offline problem to optimality, has no guarantee of meeting the best possible SCR_{Γ} .

4.4 Extension to the sequential setup-based CR analysis

The fact that the simple SFB algorithm, which does not use any information released after the beginning of the campaign, meets the tight upper bound of $\max\{0, \frac{1-GE}{1+GE}\}$ implies that the upper bound is very conservative. In this section, we propose a performance metric that is more permissive and allows us to compare the performance of different online algorithms by incorporating the additional information which is released during the campaign.

As discussed in Section 3.1, predictions are provided over τ sequential days within the OECPP framework. This means that, in addition to the initial setup Γ , the decision maker receives additional predictive information, which accumulates incrementally day by day. Furthermore, recognizing Remark 2, our model posits that, when the reoptimization is performed on any given day λ , activity schedules will remain fixed for the upcoming k days.

So, we extend the setup-based CR analysis to a dynamic setting, where the setup is updated on each day by incorporating the predictions revealed so far and the history of previous decisions. To represent the prediction history which is available at the beginning of day $\lambda \in \Lambda$, we use $\hat{\Theta}_{\lambda} = \{\hat{\theta}_{its\lambda'} : i \in I, t \in T, s \in S, \lambda' \in \{-k+1, \dots, \lambda\}\}$, i.e., $\hat{\Theta} = \hat{\Theta}_{\tau-k}$.¹³ We also let $\hat{Z}_t$ and $\hat{V}_t$ denote the values of the previous decisions

¹³ Although we have a prediction history before each day $\lambda \in \Lambda$, we do not know the quality of these predictions. Therefore, this prediction history cannot be used to build a reliable probability distribution for the rewards of activities.

corresponding to Z_{it} and V_t , respectively, for $i \in I$ and $1 \leq t < \lambda + k$. At the beginning of day $\lambda \in \Lambda$, we define $\Gamma_\lambda = \{\Gamma \cup \{Z_{it} = \hat{Z}_{it}, V_t = \hat{V}_t\} : i \in I, 1 \leq t < \lambda + k\}$ as the union of Γ and the history of previous decisions, including the fixed decisions for the upcoming k days (Remark 2), i.e., when $\lambda = -k + 1$, the condition $1 \leq t < \lambda + k$ yields an empty set; therefore no decisions are fixed at that stage and $\Gamma_{-k+1} = \Gamma$. We formally define the *sequential setups* in the context of our problem as follows.

DEFINITION 3. We define τ sequential setups corresponding to each day. On day $\lambda \in \Lambda$, we denote the setup by $\Psi_\lambda = \{\hat{\Theta}_\lambda, \Gamma_\lambda\}$. We note that Ψ_λ is a set such that each element of it contains the predictions released up to and including day λ , i.e., $\hat{\Theta}_\lambda$, and the history of the decisions of the politician up to and including day λ , i.e., Γ_λ . We also define $\hat{\Theta}_\lambda^{next}$ as the set of predictions that are not revealed yet and will be revealed after day λ . For a given setup Ψ_λ on day $\lambda \in \Lambda$, we define the sequential setup-based CR ($SSCR_{\Psi_\lambda}$) of a policy ALG_λ as:

$$0 \leq SSCR_{\Psi_\lambda} = \inf_{\hat{\Theta}_\lambda^{next}, \Theta} \frac{Reward(\Psi_\lambda, \Theta) + \mathbb{E}[ALG_\lambda(\Psi_\lambda, \hat{\Theta}_\lambda^{next}, \Theta)]}{OPT(\Gamma, \hat{\Theta})} \leq 1,$$

where $Reward(\Psi_\lambda, \Theta)$ represents the *actual* collected reward from decisions already enforced in the previous days given Ψ_λ with respect to Θ , and $\mathbb{E}[ALG_\lambda(\Psi_\lambda, \hat{\Theta}_\lambda^{next}, \Theta)]$ denotes the (expected) *actual* collected reward by the policy ALG_λ in the future days given Ψ_λ with respect to $\hat{\Theta}_\lambda^{next}$ (which is input to ALG_λ day-by-day and affects its decisions) and Θ (which is unknown to ALG_λ but used to compute its actual total reward). Furthermore, $OPT(\Gamma, \hat{\Theta})$ in the denominator represents an upper bound for the offline optimum presented in Lemma EC.1. We recall that using this upper bound is standard practice (see Remark 3).

$SSCR_{\Psi_\lambda}$ in Definition 3 provides a guarantee on the expected objective function value of any online algorithm that starts from the beginning of the campaign, achieves the setup Ψ_λ on day λ , and applies the policy ALG_λ from day λ up to the end of the campaign, against the upper bound on the offline optimum obtained with full access to $\hat{\Theta}$ from the beginning of the campaign.

In essence, $SSCR_{\Psi_\lambda}$ restricts the analysis to the family of problem instances corresponding to the observable problem inputs at the beginning of day λ , along with the fixed decisions made before day λ . From this perspective, any online algorithm for solving the OECPP can be viewed as a policy, denoted by ALG , which determines the future schedule of the politician at the beginning of each day $\lambda \in \Lambda$ based on $\hat{\Theta}_\lambda$ and Γ_λ . To operationalize $SSCR_{\Psi_\lambda}$ and compare reoptimization policies, we next present a reoptimization formulation that updates the remaining schedule given the prediction history and the decisions fixed within the freeze window.

4.5 A reoptimization framework

We note that on day $\lambda = -k + 1$, the mathematical formulation in Section 3.3 is solved using the initial predictions to determine the schedule for the first day of the campaign. Given the setup on day $\lambda \in \{-k +$

$2, \dots, \tau - k\}$, our reoptimization formulation fixes decisions through day $\lambda + k - 1$ and optimizes the schedule from day $\lambda + k$ onward. This formulation uses an arbitrary sample of predictions (chosen by the decision-maker) for $i \in I$, $t \in T$, and $s \in S$ from $\hat{\Theta}_\lambda$; that is, by day λ there are λ released predictions for each (i, t, s) , and a sample corresponds to selecting one of these predictions (or a linear combination of some of them). We denote this arbitrary prediction sample by $\hat{\Theta}_\lambda^{arb} = \{\hat{\theta}_{its\lambda}^{arb} : i \in I, t \in T, s \in S\}$. Assuming the model is reoptimized on day λ , the fixed values of decision variables Z_{it} and V_t denoted by $\hat{Z}_{it}$ and $\hat{V}_t$ are taken as input parameters for $1 \leq t < \lambda + k$, and the model optimizes the schedule from day $\lambda + k$ onward. An illustrative example for the reoptimization framework is presented in E-Companion EC.2.

Consequently, we modify the model in (1)-(10) by updating the objective function to incorporate $\hat{\Theta}_\lambda^{arb} = \{\hat{\theta}_{its\lambda}^{arb} : i \in I, t \in T, s \in S\}$.

$$\max \sum_{s \in S} \sum_{i \in I} \sum_{t \in T} X_{its} \hat{\theta}_{its\lambda}^{arb} \quad (11)$$

We also add the variable fixation constraints (12) - (13) to the model in (1)-(10).

$$Z_{it} = \hat{Z}_{it}, \quad i \in I, 1 \leq t < \lambda + k \quad (12)$$

$$V_t = \hat{V}_t, \quad 1 \leq t < \lambda + k \quad (13)$$

Building on this reoptimization framework, we now formulate a deterministic policy that leverages the available predictions on each day to achieve the best possible sequential setup-based CR on any given setup.

4.5.1 A deterministic algorithm with the best possible $SSCR_{\Psi_\lambda}$ on any Ψ_λ

We first define a deterministic policy that achieves the best possible $SSCR_{\Psi_\lambda}$ for any Ψ_λ . Then, we discuss the value of the sequential setup-based CR analysis.

DEFINITION 4. (IFBMAX Algorithm:) On day $\lambda \in \Lambda$, for $i \in I$, $t \in T$, and $s \in S$, let $\hat{\theta}_{its\lambda}^{max} = \max\{\hat{\theta}_{its(-k+1)}, \dots, \hat{\theta}_{its\lambda}\}$ and $\hat{\Theta}_\lambda^{max} = \{\hat{\theta}_{its\lambda}^{max} : i \in I, t \in T, s \in S\}$. Consider an arbitrary setup Ψ_λ and let ALG_λ^{max} be the solution of the reoptimization formulation given in Section 4.5 with respect to $\hat{\Theta}_\lambda^{max}$. Hereafter, we refer to the deterministic online algorithm which on day $\lambda \in \Lambda$ operates based on ALG_λ^{max} as the *iterative formulation-based algorithm with maximum predictions* (IFBMAX).

Based on Definition 4, when new predictions become available on day $\lambda \in \Lambda$, the IFBMAX algorithm takes into account $\hat{\Theta}_\lambda^{max}$ to update the decisions. This is why we refer to it as the IFBMAX algorithm.

PROPOSITION 2. For $\lambda \in \Lambda$, given any Ψ_λ , the best possible policy on day λ is ALG_λ^{max} from an $SSCR_{\Psi_\lambda}$ perspective.

Proposition 2 is significant because it determines the best worst-case policy for each day based on the available information, which accumulates sequentially starting from the first day of the campaign. Next, we discuss the value of reoptimization by ALG_λ^{max} on days $\lambda \in \Lambda$.

COROLLARY 1. *Let $OBJ(ALG_\lambda^{max})$ denote the total reward collected by ALG_λ^{max} (Definition 4) given Ψ_λ with respect to $\hat{\Theta}_\lambda^{max}$ during the campaign, i.e., including the reward from the fixed decisions already made in previous days, as well as the reward of the activities in the upcoming days, which have not yet become fixed. The IFBMAX algorithm achieves a better sequential setup-based CR compared to the upper bound in Theorem 1 on a subset of OECPP instances in which for at least one $\lambda < \tau - k$, it holds that $OBJ(ALG_\lambda^{max}) < OBJ(ALG_{\lambda+1}^{max})$. Furthermore, the IFBMAX algorithm maintains an SCR_Γ at least equal to the upper bound in Theorem 1 on the remaining OECPP instances.*

We recall that a deterministic algorithm which consistently uses ALG_λ^{max} on days $\lambda \in \Lambda$ is equivalent to the IFBMAX algorithm, i.e., pursuing the IFBMAX algorithm on each day is the best possible policy from an $SSCR_{\Psi_\lambda}$ perspective (Proposition 2). Consequently, Corollary 1 implies that while the IFBMAX algorithm maintains the optimality guarantee from the SCR_Γ perspective in worst-case scenarios for $\hat{\Theta}$, it guarantees the sequential setup-based CR improvement through reoptimization where applicable in non-worst-case scenarios for $\hat{\Theta}$.

4.6 Reoptimization with randomization

Having identified the best deterministic reoptimization policy, we now extend the framework to randomized policies that leverage multiple prediction samples at each reoptimization to enhance robustness and operational performance. According to the extant literature, randomization can offer improved average-case performance, robustness against adversarial inputs, and potential gains in practical implementations of online algorithms (Ma and Simchi-Levi 2020, Niazadeh et al. 2023, Ma et al. 2021, Wang et al. 2022). In this section, we propose a family of randomized algorithms that achieve a performance guarantee from an $SSCR_{\Psi_\lambda}$ viewpoint for any Ψ_λ , which is a worst-case measure. In Section 5, we will compare these algorithms based on their average-case performance in simulation scenarios.

Contrary to deterministic algorithms, where the policy ALG yields a deterministic output, in the case of randomized algorithms, the policy is defined by (1) a set of deterministic policies and (2) a probability distribution over that set, according to which one policy is selected and applied on each day. In the following, we introduce a family of randomized algorithms with $SSCR_{\Psi_\lambda}$ guarantee on any Ψ_λ .

Our randomized policy is constructed as follows. At the beginning of day $\lambda \in \{-k + 2, \dots, \tau - k\}$, we take a non-empty set of prediction samples, namely χ , from $\hat{\Theta}_\lambda$, i.e., $\chi = \{\hat{\Theta}_\lambda^{arb_1}, \hat{\Theta}_\lambda^{arb_2}, \dots, \hat{\Theta}_\lambda^{arb_{|\chi|}}\}$ (see $\hat{\Theta}_\lambda^{arb}$ in Section 4.5). We remark that inevitably $|\chi| = 1$ on day $\lambda = -k + 1$ due to the existence of only one prediction sample. However, $|\chi|$ can be decided by the decision maker in the other days. We include a sample corresponding to $\hat{\Theta}_\lambda^{max}$ in χ to ensure the sequential setup-based CR (Definition 3) guarantee of our algorithm. Then, we apply the reoptimization formulation in Section 4.5 and obtain the optimal schedule corresponding to each sample, i.e., $Schedule_1, Schedule_2, \dots, Schedule_{|\chi|}$. Once the schedules are obtained,

we compute a mapping that represents the *mapped* objective function value for each schedule with respect to $\hat{\Theta}_\lambda^{\max}$. The mapped objective function value of a schedule is the objective value obtained when that schedule is evaluated using the maximal observed prediction matrix $\hat{\Theta}_\lambda^{\max}$, regardless of which prediction sample generated it. Once the mapped objective function values are computed, we keep the schedules whose mapped objective function values are at least ρ times the objective function value of the schedule with the highest mapped objective function value (corresponding to the IFBMAX algorithm). We refer to ρ as the randomization cut-off such that $0 < \rho \leq 1$ and is specified by the decision maker. Suppose that μ schedules are selected and let $OBJ_1, OBJ_2, \dots, OBJ_\mu$ denote the mapped objective function values of them. We construct a probability distribution which allocates a selection probability of $\frac{OBJ_i}{\sum_{j=1}^{\mu} OBJ_j}$ to the solution corresponding to OBJ_i for $i = 1, 2, \dots, \mu$. Finally, we choose a schedule based on the described probability distribution. We refer to the algorithm that applies this policy as the *randomized iterative formulation-based algorithm* (RIFB), which is presented in Algorithm 1.

REMARK 4. (Randomization as a decision lever:) In our model, the randomization cut-off ρ is a parameter chosen by the campaign manager. Rather than fixing ρ , in our numerical analysis, we tune it via simulation (Section 5.4) to identify the performance-maximizing level. In practice, the primary motivation for randomization is to support decision-making in situations where multiple high-quality solutions exist under the available information, yet there is no clear basis for selecting one over another. Supporting this rationale, evidence from social science literature (Ong and Qiu 2023) indicates that individuals often value the ability to randomize when facing conflicting or incomplete preferences, and that doing so can lead to improved decision utility.

Algorithm 1 Randomized Iterative Formulation-Based Algorithm (RIFB)

```

1: Initialization ( $\lambda = -k + 1$ ):
2:   Apply the formulation in Section 3.3 to the initial prediction sample
3:   Implement the obtained schedule for the first day of the campaign
4: for each  $\lambda \in \{-k + 2, \dots, \tau - k\}$  do
5:   for  $i = 1$  to  $|\chi|$  do
6:     Sample predictions from  $\hat{\Theta}_\lambda$ 
7:     Obtain  $Schedule_i$  using reoptimization formulation in Section 4.5
8:     Compute mapped objective values for  $Schedule_i$  with respect to  $\hat{\Theta}_\lambda^{\max}$ 
9:   end for
10:  Select schedules with mapped objectives at least  $\rho$  times the maximum mapped objective
11:  Record the number of selected schedules as  $\mu$ 
12:  Label mapped objectives of selected schedules by  $OBJ_1, \dots, OBJ_\mu$ 
13:  Construct probability distribution based on  $OBJ_1, OBJ_2, \dots, OBJ_\mu$ 
14:  Choose a schedule based on this probability distribution
15: end for

```

In Proposition 3, we formally establish the performance guarantee of the RIFB algorithm. Specifically, we show that:

PROPOSITION 3. *For any Γ and Ψ_λ , and for any χ with $|\chi| \geq 1$, the RIFB algorithm achieves an SCR_Γ and $SSCR_{\Psi_\lambda}$ of at least $\rho^{\tau-1}$ times the corresponding guarantees of the IFBMAX algorithm, where $0 < \rho \leq 1$ is the randomization cut-off.*

The following corollary follows by setting $\rho = 1$ in Proposition 3:

COROLLARY 2. *For any Γ and Ψ_λ , and for any χ with $|\chi| \geq 1$, the RIFB algorithm with $\rho = 1$, achieves the best possible SCR_Γ and $SSCR_{\Psi_\lambda}$.*

In Corollary 2, we clarify that even with $\rho = 1$, there might be instances on which applying the RIFB algorithm with $|\chi| > 1$ results in more than one schedule for the politician with the same objective function value, hence causing randomization in decision-making. When ρ is less than 1, the worst-case performance of the randomized algorithm may deteriorate but the average-case performance may improve significantly. In Section 5, we compare various deterministic algorithms including the IFBMAX algorithm, as well as various versions of the RIFB algorithm by conducting sensitivity analysis on different configurations for χ and different values for ρ to understand the average-case performance of these algorithms.

4.6.1 An illustrative example

We close Section 4 with a two-day toy example illustrating the model dynamics and the associated performance guarantees.

Table 2 An illustrative example to showcase Θ , $\hat{\Theta}$, and their components

Reward Matrices			Activity Day	State A	State B
$\hat{\Theta}$	$\hat{\Theta}_1$	$\hat{\Theta}_1^1$	1	$\hat{\Theta}_{A111} = 105$	$\hat{\Theta}_{B111} = 105$
			1	$\hat{\Theta}_{A121} = 52.5$	$\hat{\Theta}_{B121} = 52.5$
			2	$\hat{\Theta}_{A211} = 2002$	$\hat{\Theta}_{B211} = 1500$
			2	$\hat{\Theta}_{A221} = 1001$	$\hat{\Theta}_{B221} = 750$
	$\hat{\Theta}_2$	$\hat{\Theta}_1^1$	1	$\hat{\Theta}_{A111} = 105$	$\hat{\Theta}_{B111} = 105$
			1	$\hat{\Theta}_{A121} = 52.5$	$\hat{\Theta}_{B121} = 52.5$
			2	$\hat{\Theta}_{A211} = 2002$	$\hat{\Theta}_{B211} = 1500$
			2	$\hat{\Theta}_{A221} = 1001$	$\hat{\Theta}_{B221} = 750$
		$\hat{\Theta}_2^2$	1	$\hat{\Theta}_{A112} = 105$	$\hat{\Theta}_{B112} = 105$
			1	$\hat{\Theta}_{A122} = 52.5$	$\hat{\Theta}_{B122} = 52.5$
			2	$\hat{\Theta}_{A212} = 1000$	$\hat{\Theta}_{B212} = 5992$
			2	$\hat{\Theta}_{A222} = 500$	$\hat{\Theta}_{B222} = 2996$
Θ			1	$\Theta_{A11} = 420$	$\Theta_{B11} = 60$
			1	$\Theta_{A12} = 210$	$\Theta_{B12} = 30$
			2	$\Theta_{A21} = 1144$	$\Theta_{B21} = 3424$
			2	$\Theta_{A22} = 572$	$\Theta_{B22} = 1712$

Key parameters: We first present the key parameters in our example.

- Γ : We consider a simple example with two states: A and B, and a two-day campaign. In this example, we assume that $\alpha = 1$, $\beta = 0$, and $\bar{\beta} = 2$. For simplicity, we also assume that k is zero, i.e., the politician can make decisions for the same day at the beginning of that day.
- Ψ_λ : Since $k = 0$, we have $\Psi_1 = \{\hat{\Theta}_1, \Gamma_1\} = \{\hat{\Theta}_1, \Gamma\}$. On day 2, the setup depends on the realized decision on day 1. Specifically, if no activity is conducted on day 1, then $\Psi_2 = \{\hat{\Theta}_2, \Gamma \cup \{Z_{A1} = 0, Z_{B1} = 0, V_1 = 1\}\}$; if an activity is conducted in state A on day 1, then $\Psi_2 = \{\hat{\Theta}_2, \Gamma \cup \{Z_{A1} = 1, Z_{B1} = 0, V_1 = 0\}\}$; and if an activity is conducted in state B on day 1, then $\Psi_2 = \{\hat{\Theta}_2, \Gamma \cup \{Z_{A1} = 0, Z_{B1} = 1, V_1 = 0\}\}$.
- GE: We fix GE to 0.75 in this example. We highlight that GE is not known to the decision-maker.
- Θ : We provide a random realization of Θ in Table 2 which respects $GE = 0.75$. We emphasize that Θ will never be known to the decision-maker.
- $\hat{\Theta}$: The prediction matrix (shown by $\hat{\Theta}$ in Table 2) is revealed to the decision-maker day-by-day. On day 1, only $\hat{\Theta}_1$ is available to the decision-maker. On day 2, both $\hat{\Theta}_1$ and $\hat{\Theta}_2$ are available to the decision-maker, i.e., $\hat{\Theta}_1 = \{\hat{\Theta}_1^1\}$ and $\hat{\Theta}_2 = \{\hat{\Theta}_1^1, \hat{\Theta}_2^2\}$, where the superscripts are used to avoid self-reference in the sets $\hat{\Theta}_1$ and $\hat{\Theta}_2$.

Computation of the upper bound on the offline optimum, $OPT(\Gamma, \hat{\Theta})$: According to Lemma EC.1, the upper bound on the offline optimum can be computed by multiplying $\frac{1}{1-GE}$ and the best achievable objective function value with respect to the *minimum* predictions for each activity. In this example, GE is 0.75 (see Definition 2) and the best achievable objective function value considering minimum predictions for each activity is $105 + 1500 = 1605$ corresponding to conducting an activity in state A on day 1, and conducting an activity in state B on day 2 (see Table 2). Hence, the upper bound on the offline optimum can be calculated as $\frac{1}{1-0.75} \times 1605 = 6420$. This means in our example, there is no scenario for Θ in which the offline optimum can collect a total reward of higher than 6420.

An implementation of the randomized algorithm: We now present an example of the implementation of the RIFB algorithm with $\chi = \{\hat{\Theta}_2^{max}, \hat{\Theta}_1\}$ and $\rho = 0.9$ as follows. On day 1, there is no randomization because χ is a singleton set. Hence, the algorithm conducts an activity in state B on day 1 (see line 3 in Algorithm 1). On day 2, the algorithm obtains two different schedules based on the two prediction samples (see line 7 in Algorithm 1), i.e., $Schedule_1$ and $Schedule_2$ corresponding to $\hat{\Theta}_2^{max}$ and $\hat{\Theta}_1$, respectively.¹⁴ Then, the algorithm finds the mapped objective function values of these two schedules based on maximal observed predictions, i.e., these mapped objective function values for $Schedule_1$ and $Schedule_2$ are $52.5 + 2996 = 3048.5$ and $105 + 2002 = 2107$, respectively (see line 8 in Algorithm 1). Then, since $\rho = 0.9$, only $Schedule_1$ is selected for decision-making, i.e., $\mu = 1$ in line 11 of Algorithm 1. Hence, the performance of the algorithm in this example matches that of the IFBMAX algorithm. It should be noted that randomization will not be activated in this example unless ρ is reduced to approximately 0.691 or lower.

¹⁴ $Schedule_1$ corresponds to conducting an activity in state B on both Day 1 and Day 2, whereas $Schedule_2$ corresponds to conducting an activity in state B on Day 1 and in state A on Day 2.

Performance of the algorithm: Since $\hat{\Theta}$ is non-worst-case in this example, $SSCR_{\Psi_2}$ of the above-discussed randomized algorithm corresponding to $Schedule_1$, which is $\frac{30+1712}{6420} = \frac{1742}{6420} \approx 0.271$ (Corollary 1), surpasses the tight upper bound on SCR_{Γ} for $GE = 0.75$, which is $\frac{1-0.75}{1+0.75} \approx 0.142$ (Theorem 1).

5 Empirical analysis

In this section, we evaluate the algorithms proposed in Section 4 from an average-case perspective using the empirical CR metric. We first define the metric, then describe the U.S. presidential election dataset and the prediction-generation process, and finally compare deterministic and randomized policies, including a counterfactual analysis of the realized 2024 campaign schedule.

5.1 Performance criterion

The analysis in Section 4 is based on a theoretical CR perspective which is of a *worst-case* nature. Although such analysis constructs a theoretical foundation for analyzing the OECPP, it does not comment on the *average* performance of online algorithms. This motivates us to investigate the following questions. How do the alternative online algorithms discussed in Section 4 perform from an average-case perspective (see Ma et al. (2021))? What is a suitable metric to analyze the average performance of online algorithms?

Accordingly, we analyze the performance of the online algorithms empirically by considering instances that simulate the real-world characteristics of the OECPP in the context of U.S. presidential elections. The performance criterion used to evaluate the average-case performance of online algorithms is *empirical CR* (ECR), which is defined as follows.

DEFINITION 5. For a given subset of instances Δ , where each instance is defined as a tuple $(\Gamma, \hat{\Theta}, \Theta)$ consisting of a setup (Γ) , a global reward prediction matrix $(\hat{\Theta})$ that is revealed to the online algorithm day-by-day, and a global actual reward matrix (Θ) which is unknown to the online algorithm, the ECR of an online (deterministic or randomized) algorithm is the ratio of the total objective value obtained by the online algorithm to the total objective value of the offline optimum (*OFF*) across those instances (Wang et al. 2022, Shiri et al. 2024), i.e.,

$$ECR = \frac{\sum_{(\Gamma, \hat{\Theta}, \Theta) \in \Delta} \mathbb{E}[ALG(\Gamma, \hat{\Theta}, \Theta)]}{\sum_{(\Gamma, \hat{\Theta}, \Theta) \in \Delta} OFF(\Gamma, \Theta)},$$

where $ALG(\Gamma, \hat{\Theta}, \Theta)$ represents the (expected) *actual* objective function value of the online algorithm with respect to Γ , $\hat{\Theta}$, and Θ . Furthermore, $OFF(\Gamma, \Theta)$ denotes the offline optimum, computed with full knowledge of Γ and Θ from the outset.

We clarify that in the ECR, the performance of the online algorithm on each tested instance is compared with the exact offline optimum value computed for that instance with access to Θ . That is, contrary to Definitions 1 and 3, where the upper bound on the offline optimum is considered in the denominator of the

setup-based CR due to the difficulty of incorporating the exact offline optimum in the theoretical analysis, in the ECR we use the exact offline optimum in the denominator. This leads to a more accurate evaluation of the average-case performance of alternative online algorithms.

5.2 Dataset

We conducted our empirical experiments on a dataset constructed from available data on previous U.S. presidential elections, supplemented with assumptions grounded in the political science literature (Section 2.2). To generate each problem instance, we construct three components: (1) the setup of deterministic input parameters (Γ), (2) the actual reward matrix (Θ), and (3) the reward prediction matrix ($\hat{\Theta}$). Next, we describe the problem instances in our experiments and explain how they were generated.

5.2.1 Generation of Γ

Because our case study focuses on the U.S. presidential election, we fix Γ in all experiments. That is, we consider all 50 U.S. states plus Washington, D.C. (i.e., $|I| = 51$). Furthermore, as the U.S. election campaign typically starts after each party's convention and spans 50–70 days, we consider a campaign of 65 days in our experiments ($|T| = 65$) with $k = 5$, so that the first optimization occurs 5 days before the start of the campaign and 70 days before Election Day. In prior U.S. presidential elections, most campaign days were confined to a single state; even when several events occurred on the same day, they typically took place within that state. To reflect this pattern and for tractability, we set $\alpha = 1$ in our experiments. Moreover, historical data indicate that candidates typically include an average of 3 rest days during the initial weeks of the campaign and no rest days in the final two weeks; therefore, we set $\beta = 3$ in our experiments. Similarly, historical data show that candidates rarely conduct activities in the same state on more than 2 days within any consecutive seven-day period. Hence, we set $\bar{\beta} = 2$ to enforce this limit.¹⁵

5.2.2 Generation of Θ

As discussed in Section 3.1, the elements in Θ are characterized as $\theta_{its} = \theta_i \times \Phi_t \times \Xi_s$ for $i \in I$, $t \in T$, $s \in S$, and are unknown to the decision maker. It is difficult to accurately calculate Θ because it remains unknown in practice, even when the election outcome is known. In this paper, for benchmarking purposes, we compute Θ based on the main variables influencing the rewards and insights from the political science literature (Section 2.2). For modeling θ_i , we incorporate four key state-specific factors: Electoral College Votes (EC_i), Criticality Factor (CF_i), Resistance Factor (RF_i), and Swing Indicator Function (SS_i).

¹⁵ Based on historical data from U.S. elections prior to 2024—which show that no state hosted main campaign events (e.g., rallies or candidate speeches) on more than 12 days during the last 65 days of the campaign—we restrict the domain of the set S to values less than or equal to 12 in our implementation of the models in Sections 3.3 and 4.5, in order to reduce computational run time.

EC_i quantifies the relative weight of a state's *electoral college votes*.¹⁶ CF_i measures how sensitive a state's election outcome is to shifts in voter preference, based on the *most recent* election—specifically, the 2020 election in our case study. RF_i captures a state's *historical* tendency to switch its support between political parties, based on election data from 2000 onward. SS_i measures the relative importance of a state as a battleground in the *current* election, based on the 2024 election in our case study.¹⁷ We refer the reader to E-Companion EC.4 for details about the calculation of θ_i in our experiments.

We specify Φ_t as a U-shaped function of t , and Ξ_s as a decreasing function of s , guided by the literature in Section 2.2. Specifically, we set $\Phi_t = (t - \frac{1+\tau}{2})^2 + C$, where the constant additive term C is chosen to ensure that the minimum value of Φ_t remains a specified fraction of its endpoint values. In our implementation, we set $C = (\frac{\tau-1}{4})^2$, which corresponds to a moderate level of temporal smoothing and can be adjusted according to empirical calibration or design preference. Furthermore, to calibrate Ξ_s to match the historical upper limit on the number of days with activities conducted in the same state during U.S. presidential elections (no more than 12 in the final 65 days leading up to election day), we set $\Xi_s = 1 - 0.04 \times (s - 1)$, so that a 4% penalty is applied for each additional activity in that state over the campaign. The choice of the 4% factor was determined through empirical tuning to ensure that the resulting schedules respect this upper limit.

5.2.3 Generation of $\hat{\Theta}$

We analyze the performance of our online algorithms with respect to GE in both our theoretical and empirical analyses. Accordingly, we generate $\hat{\Theta}$ based on a truncated multivariate normal distribution, constructed using GE, a mean vector (Θ) and a covariance matrix that captures the correlations in rewards (by state) based on historical U.S. presidential election data.¹⁸ To generate predictions at the desired GE levels, we sample from this distribution and truncate the samples to enforce GE bounds of 10%, 30%, 50%, 70%, and 90%. Specifically, for each GE level, we take 10 samples—resulting in a total of $5 \times 10 = 50$ scenarios for $\hat{\Theta}$. Next, we incorporate unforeseen shocks into our prediction framework to reflect real-world scenarios more accurately.¹⁹ In our planning horizon, we assume that one shock is encountered per week, i.e., for each generated prediction seed, we update the predictions within their respective error intervals every seven days.

¹⁶ While EC_i captures a state's electoral weight, it alone does not fully reflect the state's realized importance. States with large electoral vote counts may consistently favor one political party, reducing their relevance as battlegrounds. For instance, Texas—with 40 electoral votes—has reliably supported the Republican Party in recent decades, as evidenced by its voting record and polling data from recent election cycles. To address this, we also incorporate the factors CF_i , RF_i , and SS_i , which measure how sensitive a state's election outcome is to shifts in voter preference, based on publicly available data.

¹⁷ We use pre-campaign poll data from 2024 to construct this measure.

¹⁸ Recall that the reliability of these predictions is unknown to the decision-maker.

¹⁹ These shocks were incorporated into our predictions due to their frequent occurrence during election campaigns. For example, in late October 2016, the FBI unexpectedly announced the reopening of its investigation into Hillary Clinton's emails. This national-level shock occurred shortly before the election and coincided with a noticeable tightening of polling margins in key battleground states such as Michigan and Pennsylvania. Consequently, the perceived value of campaigning in those states changed significantly. See <https://time.com/4560064/fbi-hillary-clinton-decision-what-to-know/>.

5.3 Tested algorithms

We evaluate the following algorithms in our empirical analysis.

Randomized algorithms. We test different versions of the RIFB algorithm as follows. We consider four different options for the prediction samples (i.e., $\hat{\Theta}_\lambda^{arb}$) in the RIFB algorithm, which correspond to the most intuitive policies.

1. $\hat{\Theta}_1 = \{\hat{\theta}_{its(-k+1)} : i \in I, t \in T, s \in S\}$ which represents the initial predictions revealed on day $\lambda = -k + 1$,
2. $\hat{\Theta}_\lambda^{max} = \{\hat{\theta}_{its\lambda}^{max} : i \in I, t \in T, s \in S\}$ which represents the maximum realized predictions (see Definition 4),
3. $\hat{\Theta}_\lambda^{avg} = \{\hat{\theta}_{its\lambda}^{avg} : i \in I, t \in T, s \in S\}$ which represents the average predictions, where $\hat{\theta}_{its\lambda}^{avg}$ is the average of released values, and
4. $\hat{\Theta}_\lambda^\lambda = \{\hat{\theta}_{its\lambda} : i \in I, t \in T, s \in S\}$ which represents the most recent predictions.

We consider $|\chi| \in \{1, 2, 4\}$ and $\rho \in \{0.7, 0.8, 0.9, 1\}$. For brevity, we omit $|\chi| = 3$, which yields similar conclusions. For $|\chi| = 1$, the only prediction sample is $\hat{\Theta}_\lambda^{max}$, corresponding to the IFBMAX algorithm, which is optimal from a worst-case perspective. For $|\chi| = 2$, we consider three combinations: $(\hat{\Theta}_\lambda^{max}, \hat{\Theta}_1)$, $(\hat{\Theta}_\lambda^{max}, \hat{\Theta}_\lambda^{avg})$, and $(\hat{\Theta}_\lambda^{max}, \hat{\Theta}_\lambda^\lambda)$. For $|\chi| = 4$, we include all four prediction samples. Therefore, we consider five scenarios for χ and four values of ρ , yielding $5 \times 4 = 20$ variants of the RIFB algorithm.

Deterministic algorithms. We also test four different deterministic algorithms, including the SFB algorithm and the IFBMAX algorithm. The other two algorithms are iterative formulation-based algorithms that use $\hat{\Theta}_\lambda^{avg}$ and $\hat{\Theta}_\lambda^\lambda$ to update their solution at their reoptimization points, and we refer to them as IFBAVG and IFB $_\lambda$, respectively.

5.4 Identifying the best randomized algorithm

The performance of randomized algorithms depends on two key parameters: (1) the set of prediction samples used for randomization χ , and (2) the randomization cut-off ρ . In this section, we propose a simulation-based framework to determine the best values for the randomization parameters among 20 configurations corresponding to different versions of the RIFB algorithm described in Section 5.3, for our case study. We recall from Section 4.1 that randomized algorithms may produce different outcomes in different executions on the same instance. To address this, each variant of the RIFB algorithm was executed 10 times across each of the 50 scenarios for $\hat{\Theta}$ described in Section 5.2.3, and its average performance across these 10 executions was considered. Also, recall from Section 5.3 that we test 20 versions of the RIFB algorithm. This yields a total of $20 \times 10 \times 50 = 10,000$ experiments for the randomized algorithms.

While our ILP formulation, given in Section 3.3, is capable of finding the optimal solution for instances with 65 days in a reasonable time (less than 2 hours), it is impractical to perform 10,000 experiments considering a 65-day campaign period. Therefore, in our experiments for tuning the randomization parameters, we

consider a shortened campaign horizon of 35 days.²⁰ Using the tuned parameters from these experiments, we run a full-scale counterfactual analysis considering the full 65-day campaign in Section 5.6.

In Table 3, we report the ECR performance of our randomized algorithm for all reasonable combinations of χ and ρ where the set of tested instances, Δ in Definition 5, corresponds to the 50 realizations of $\hat{\Theta}$ explained in Section 5.2.3. We consistently include $\hat{\Theta}_\lambda^{\max}$ in the set of prediction samples χ to ensure a worst-case CR guarantee. Note that the cases with $\chi = \{\hat{\Theta}_\lambda^{\max}\}$ represent the non-randomized IFBMAX algorithm. We observe that (1) randomization can improve ECR performance, (2) adding more prediction samples to the set χ does not necessarily result in better ECR, and (3) the cut-off parameter ρ impacts ECR performance in a non-monotone manner; that is, ECR first increases and then decreases with ρ . Overall, Table 3 shows that the optimal combination of parameters resulting in the highest ECR is given by $\chi = \{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$, and $\rho = 0.9$.

Two key insights emerge from this analysis: On the one hand, by including more prediction samples in the randomization set besides $\hat{\Theta}_\lambda^{\max}$, one can improve the average-case performance by spreading risk across alternative solutions of similar quality, thus enhancing overall performance and reducing the impact of any single adverse outcome. On the other hand, including too many prediction samples in the set χ , along with a low randomization cut-off, can lead to decisions driven by lower-quality prediction samples within the set χ , ultimately reducing the ECR.

Table 3 Average ECR performance of randomized algorithms for various scenarios for χ and different values for ρ

χ	$\rho = 1$	$\rho = 0.9$	$\rho = 0.8$	$\rho = 0.7$
$\{\hat{\Theta}_\lambda^{\max}\}$	92.66%	92.66%	92.66%	92.66%
$\{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$	92.66%	94.28%	91.50%	90.58%
$\{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_1\}$	92.66%	90.96%	89.24%	88.32%
$\{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^\lambda\}$	92.66%	91.46%	90.08%	89.26%
$\{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}, \hat{\Theta}_1, \hat{\Theta}_\lambda^\lambda\}$	92.66%	92.08%	90.46%	89.84%

Figure 2 illustrates the performance of multiple randomized algorithms corresponding to different scenarios for χ , with ρ fixed at 0.9, in terms of ECR across various GE levels. The results suggest that the randomized algorithm which uses $\chi = \{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$ on day $\lambda \in \Lambda$, performs better across all GE levels. When GE is low, randomization offers limited benefit: the deterministic IFBMAX algorithm (with $\chi = \{\hat{\Theta}_\lambda^{\max}\}$) performs nearly identically to the best-performing randomized algorithm (with $\chi = \{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$). As GE increases, however, the performance gap between IFBMAX and the best-performing randomized algorithm widens rapidly, i.e., the benefit of randomization increases. We also note that ECR performance degrades with GE level for all algorithms.

²⁰ While 35 days is sufficient to understand the behavior of the RIFB algorithm under different parameter settings, in practice, political parties can conduct similar simulation-based tuning over the full campaign period by allocating more computational resources.

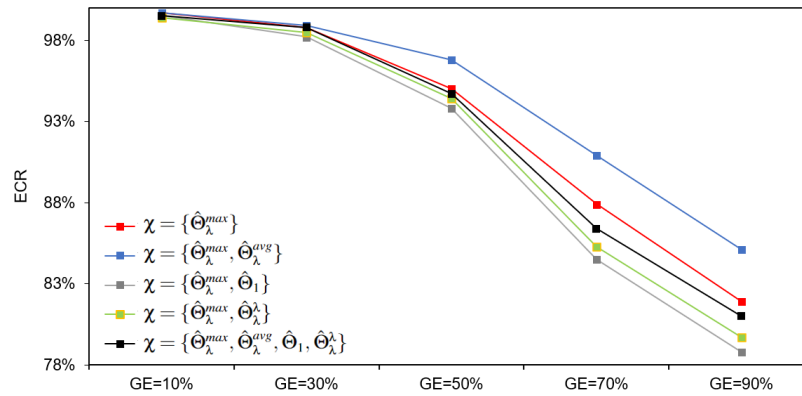


Figure 2 Performance of the RIFB algorithms with $\rho = 0.9$ across five different scenarios for χ . The x-axis represents the error level, while the y-axis shows the performance of the algorithms in terms of ECR. Each curve corresponds to a different algorithm with a specific χ .

Figure 3 depicts the ECR values of the randomized RIFB algorithm with $\chi = \{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$ for varying levels of ρ and GE. The results suggest that for low GE values, randomization has limited effect on the ECR of the algorithm, and its impact increases as GE rises. For a given level of GE, we observe that randomization may improve ECR. The biggest improvement in ECR is observed when the randomization cut-off (ρ) is around 0.9. For higher and lower values of the randomization cut-off, the benefit of randomization diminishes, and it can even lead to a reduction in ECR. Thus, properly selecting the randomization cut-off is key to benefiting from randomized algorithms.

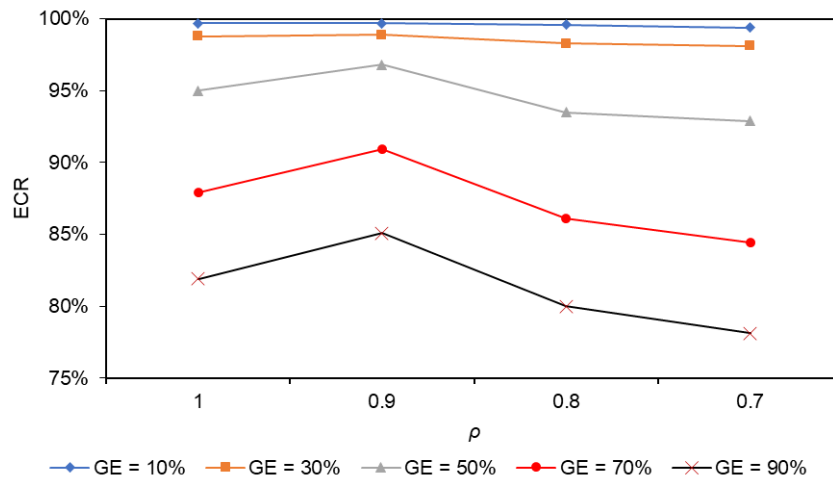


Figure 3 Performance of the randomized algorithm according to varying ρ values for $\chi = \{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$

5.5 Assessing the impact of randomization

In Figure 4, we compare the performance of the best randomized algorithm ($\chi = \{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$ and $\rho = 0.9$), which we refer to as RIFB(best), with the deterministic algorithms. The results show that the randomized

algorithm RIFB(best) outperforms the deterministic algorithms across all levels of GE, with the performance gap increasing as GE increases. In general, the ECR of the randomized algorithm is much more robust to increases in GE compared to the deterministic algorithms.

This superior performance can be attributed to the use of a cut-off parameter ρ , whose optimal value ($\rho = 0.9$) was identified through simulation-based analysis in Section 5.4. This parameter generates a pool of prediction samples associated with high-quality solutions, based on the available predictions. The pool is dynamically updated after each change in the predictions, and a sample is drawn from it for reoptimization using the formulation in Section 4.5. Importantly, randomization helps mitigate the inherent biases present in deterministic policies. For example, $\hat{\Theta}_1$ disregards new information, while $\hat{\Theta}_\lambda^\lambda$ disregards past information. Similarly, $\hat{\Theta}_\lambda^{\max}$ and $\hat{\Theta}_\lambda^{\text{avg}}$ are sensitive only to extreme or average prediction values, respectively. In contrast, the randomized algorithm leverages a combination of prediction samples selected via simulation to produce high-quality solutions, avoiding overcommitment to any single biased perspective.

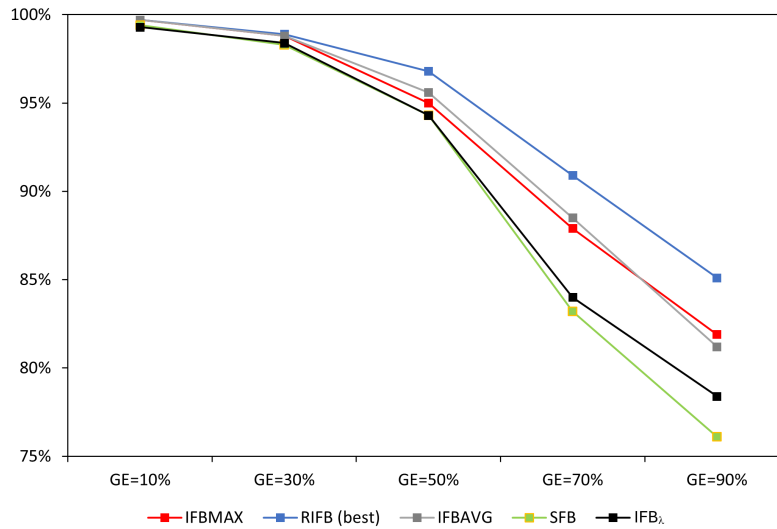


Figure 4 Comparison of the best randomized algorithm with different deterministic algorithms

5.6 Counterfactual analysis

To highlight the similarities and differences between our algorithms and real-life U.S. election campaigns, we compare the IFBMAX solution (obtained for the Democratic candidate) with the actual campaign schedule of Kamala Harris during the final 65 days leading up to the 2024 U.S. presidential election.^{21,22}

²¹ All inputs to the online algorithm, including reward predictions, are derived solely from historical election data and publicly available poll data available prior to the start of the 2024 campaign; no data from the 2024 campaign itself were used for model calibration.

²² To extract the Democratic candidate's real-life schedule, we retained, whenever relevant, only the candidate's primary daily activity to ensure comparability with our solution.

In this analysis, the IFBMAX algorithm was provided with a sample reward prediction matrix $\hat{\Theta}$ generated using the truncated multivariate normal distribution with respect to $GE = 0.5$, as described in Section 5.2.3. The comparison between the IFBMAX solution and the Democratic candidate’s actual plan is presented in Figure 5. (We also conduct a counterfactual analysis with respect to the best-performing randomized algorithm, identified in Section 5.4, in E-Companion EC.5.1, but we defer its discussion there for clarity of exposition.)

Both the IFBMAX solution and the Democratic candidate’s schedule exhibit a focus on swing states, including Pennsylvania, Michigan, Wisconsin, Georgia, Arizona, Nevada, and North Carolina. We present the detailed campaign schedule in both plans in E-Companion EC.5.2. While both plans share considerable similarities, reflecting the practical relevance of our model, they also exhibit several differences. The real-world campaign strategies are influenced by a broader set of constraints, including logistical limitations, media availability, and strategic priorities that may not be captured in our modeling framework due to the lack of publicly available data on such information. It is also possible that some differences stem from genuinely suboptimal decisions made during the actual campaign—arising from cognitive biases or strategic misjudgments—which are not uncommon in complex, high-pressure political environments (Devine 2018, Haselmayer 2019, Dommett 2019). Below, we present some of our key observations.

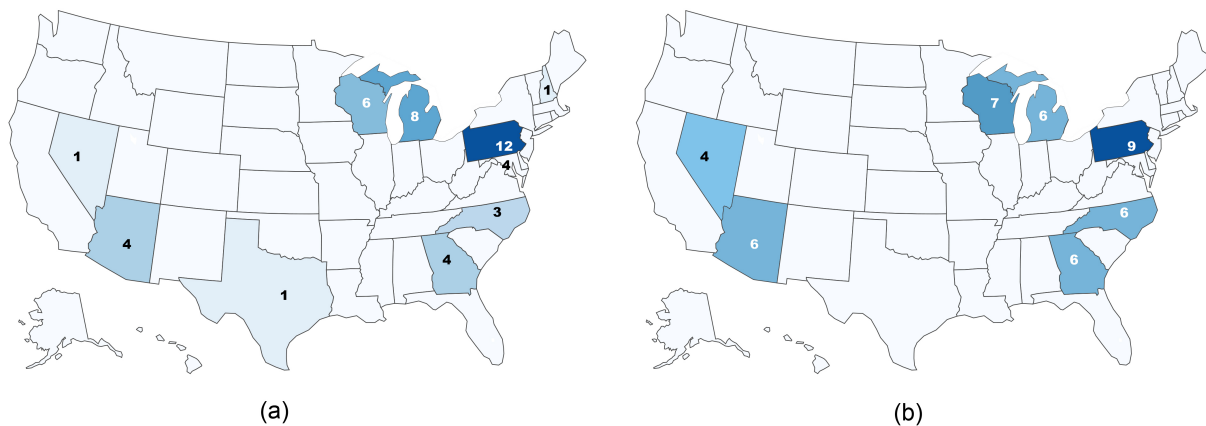


Figure 5 Comparison between (a) the Democratic candidate’s real-world plan and (b) the IFBMAX-generated schedule for the 2024 U.S. presidential election.

The IFBMAX algorithm avoids one-off campaign activities in a state and instead recommends a more consolidated effort focused on the key battlegrounds. That is, in contrast to the Democratic candidate, who visited Nevada, Texas, and New Hampshire only once, such one-off activities are avoided in the IFBMAX solution. Additionally, the Democratic candidate visited D.C. four times. While visiting D.C. offers symbolic benefits and operational convenience, critics maintain that the campaign could have gained more traction by engaging directly with voters in battleground regions through unfiltered outreach. We acknowledge that there are also political arguments supporting the Democratic candidate’s decision to spend time

in D.C. However, the IFBMAX solution focuses on key states and does not visit D.C. Finally, the IFBMAX solution directs significantly more attention to Nevada, Georgia, Arizona, and North Carolina (all of which were lost by Democrats in the 2024 election), reallocating this time either from other lost battlegrounds, such as Pennsylvania and Michigan, or from non-battleground states like Texas, New Hampshire, and D.C.

Finally, for a quantitative comparison, we evaluated the Democratic candidate’s realized 2024 campaign schedule and the IFBMAX solution under the proposed reward structure, obtaining objective function values of 493,190.39 and 657,964.87, respectively. This corresponds to a 33.41% higher objective function value, under the same benchmark, for IFBMAX relative to the Democratic candidate’s realized schedule, suggesting potential performance gains from our approach under this reward structure.

6 Conclusion

In this study, we introduce the Online Election Campaign Planning Problem, a novel adaptation of the Roaming Salesman Problem tailored for election campaign planning. The OECPP addresses the practical challenge of unknown rewards in election campaign planning by integrating real-time predictions into its methodological framework. We develop online solution methodologies that operate under sequentially received unreliable predictions and leverage the competitive ratio (CR) as a performance metric, extending the CR framework to incorporate sequentially updated information.

We draw inspiration from Ma and Simchi-Levi (2020) to analyze the performance of online algorithms based on the realized setup of input parameters. We demonstrate that the single-shot formulation-based (SFB) algorithm, which follows the optimal solution of the exact deterministic model, achieves a setup-based CR equal to the tight upper bound on the setup-based CR of online algorithms for the OECPP. The key insight is that an exact mathematical formulation applied to predictions with unknown errors is essential to guarantee the optimal policy in the worst case.

Furthermore, we extend the setup-based CR to a *sequential* setting to incorporate the effect of sequentially received unreliable predictions into the performance metric. This leads to a reoptimization framework, where the policy is dynamically updated in response to predictions over time. We show that the sequential setup-based CR of the IFBMAX algorithm, which operates based on the maximum observed predictions in each day, is greater than or equal to the tight upper bound on the setup-based CR. This implies that while the IFBMAX algorithm is robust against worst-case setups from a setup-based CR perspective, it achieves higher rewards in non-worst-case setups through reoptimization. For example, compared to the SFB algorithm, which meets the upper bound in Theorem 1, the IFBMAX algorithm collects more rewards in some setups and equal rewards in others.

Our empirical experiments, tailored to U.S. presidential elections, demonstrate that randomization can improve the empirical CR (ECR). The key factor influencing this improvement is the inclusion of additional prediction samples beyond $\hat{\Theta}_\lambda^{\max}$, which improves average-case performance by spreading risk across

multiple high-quality schedules. However, including too many prediction samples and selecting a low randomization cut-off can lead to decision-making driven by potentially inferior samples, ultimately reducing ECR.

In our U.S. presidential election-based experiments, the randomized algorithm with prediction samples $\chi = \{\hat{\Theta}_\lambda^{\max}, \hat{\Theta}_\lambda^{\text{avg}}\}$ and randomization cut-off $\rho = 0.9$ performs better across all tested GE levels. The most significant improvement in ECR is observed with a randomization cut-off around 0.9; beyond this level, the benefits diminish, reducing ECR. Therefore, selecting an appropriate randomization cut-off is crucial for maximizing the benefits of randomized algorithms. Finally, while the effectiveness of randomization increases as GE rises, it does not improve ECR at lower GE levels.

7 Limitations and future research directions

Our formulation focuses on first-order drivers of campaign scheduling under uncertain and dynamically updated reward predictions. As a result, it abstracts from some practical considerations that can affect real-life campaign planning (e.g., venue availability, security constraints, media opportunities, surrogate coordination, and candidate obligations). We therefore position OECPP as a decision-support tool that can inform planning by generating data-driven recommendations, which can then be reviewed and adjusted by campaign staff to ensure feasibility under practical constraints.

Future work could explore several directions. First, stress-testing the model under varying assumptions about the reward function could help assess its robustness. For example, future studies could analyze the sensitivity of the recommended policies to changes in the relative importance of states, timing effects, or diminishing returns, and evaluate performance under systematically biased prediction errors. Second, the model could be extended by explicitly incorporating the campaign budget as a constraint with uncertain value, updated over time as new information arrives. This would enable analysis of how both the presence of a budget constraint and uncertainty about its magnitude shape campaign strategy—particularly in local, regional, gubernatorial, and mayoral races, where budget management is a first-order concern. Third, an important direction is to adapt the model to operate at the county level by leveraging scalable heuristics and incorporating more localized data. Although computational complexity and data limitations present challenges, a county-level extension would provide finer-grained insights—especially in battleground states where changes in a few key counties can significantly impact state-level results within the Electoral College framework. Finally, extending the OECPP to settings with multiple primary candidates within a party could inform the analyses of more complex election systems, such as parliamentary elections, and help characterize how intra-party interactions affect overall strategy and outcomes.

Acknowledgments

The authors gratefully thank the reviewers of POM.

References

- Ansolabehere, Stephen, Eitan Hersh. 2012. Validation: What big data reveal about survey misreporting and the real electorate. *Political Analysis*, 20 (4), 437-459.
- Aouad, Ali, Daniela Saban. 2023. Online assortment optimization for two-sided matching platforms. *Management Science*, 69 (4), 2069-2087.
- Ausiello, Giorgio, Esteban Feuerstein, Stefano Leonardi, Leen Stougie, Maurizio Talamo. 2001. Algorithms for the on-line travelling salesman 1. *Algorithmica*, 29 (4), 560-581.
- Bartels, Larry M. 1985. Resource allocation in a presidential campaign. *The Journal of Politics*, 47 (3), 928-936.
- Beachler, Donald W, Matthew L Bergbower, Chris Cooper, David F Damore, Bas Van Dooren, Sean D Foreman, Rebecca D Gill, Henriët Hendriks, Donna Hoffmann, Rafael Jacob, et al. 2015. *Presidential Swing States: Why Only Ten Matter*. Lexington Books.
- Belenky, Alexander S. 2016. The electoral college and campaign strategies. *Who Will Be the Next President? A Guide to the US Presidential Election System*. Springer, 75-92.
- Brams, Steven J., Morton D. Davis. 1974. The 3/2's rule in presidential campaigning. *American Political Science Review*, 68 (1), 113-134.
- Broockman, David E, Joshua L Kalla. 2023. When and why are campaigns' persuasive effects small? evidence from the 2020 us presidential election. *American Journal of Political Science*, 67 (4), 833-849.
- Brox, Brian, Joseph Giammo. 2009. Late deciders in us presidential elections. *American Review of Politics*, 30 333-355.
- Campbell, W Joseph. 2024. *Lost in a Gallup: Polling failure in US presidential elections*. Univ of California Press.
- Cinar, Ahmet, F Sibel Salman, Burcin Bozkaya. 2021. Prioritized single nurse routing and scheduling for home healthcare services. *European Journal of Operational Research*, 289 (3), 867-878.
- Colantoni, Claude S., Terrence J. Levesque, Peter C. Ordeshook. 1975. Campaign resource allocations under the electoral college. *American Political Science Review*, 69 (1), 141-154.
- Devine, Christopher J. 2018. What if hillary clinton had gone to wisconsin? presidential campaign visits and vote choice in the 2016 election. *The Forum*, vol. 16. de Gruyter, 211-234.
- Dommett, Katharine. 2019. Data-driven political campaigns in practice: understanding and regulating diverse data-driven campaigns. *Internet Policy Review*, 8 (4), 1-18.
- Feng, Yiding, Rad Niazadeh. 2025. Batching and optimal multistage bipartite allocations. *Management Science*, 71 (5), 4108-4130.
- Finkel, Steven E. 1993. Reexamining the" minimal effects" model in recent presidential campaigns. *The Journal of Politics*, 55 (1), 1-21.
- Golrezaei, Negin, Hamid Nazerzadeh, Paat Rusmevichientong. 2014. Real-time optimization of personalized assortments. *Management Science*, 60 (6), 1532-1551.
- Gouleakis, Themistoklis, Konstantinos Lakis, Golnoosh Shahkarami. 2023. Learning-augmented algorithms for online tsp on the line. *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37. 11989-11996.
- Haselmayer, Martin. 2019. Negative campaigning and its consequences: a review and a look ahead. *French Politics*, 17 (3), 355-372.
- Herr, J Paul. 2002. The impact of campaign appearances in the 1996 election. *The Journal of Politics*, 64 (3), 904-913.
- Hill, Seth J., Gregory A. Huber. 2017. Representativeness and motivations of the contemporary donorate: Results from merged survey and administrative records. *Political Behavior*, 39 (1), 3-29.
- Holbrook, Thomas. 1996. *Do campaigns matter?*, vol. 1. Sage publications.
- Jaillet, Patrick, Michael R. Wagner. 2008. Generalized online routing: New competitive ratios, resource augmentation, and asymptotic analyses. *Operations Research*, 56 (3), 745-757.

- Jowell, Roger, Barry Hedges, Peter Lynn, Graham Farrant, Anthony Heath. 1993. The 1992 british election: the failure of the polls. *The Public Opinion Quarterly*, 57 (2), 238-263.
- Lykouris, Thodoris, Sergei Vassilvitskii. 2021. Competitive caching with machine learned advice. *Journal of the ACM (JACM)*, 68 (4), 1-25.
- Ma, Will, David Simchi-Levi. 2020. Algorithms for online matching, assortment, and pricing with tight weight-dependent competitive ratios. *Operations Research*, 68 (6), 1787-1803.
- Ma, Will, David Simchi-Levi, Chung-Piaw Teo. 2021. On policies for single-leg revenue management with limited demand information. *Operations Research*, 69 (1), 207-226.
- Manshadi, Vahideh, Scott Rodilitz, Daniela Saban, Akshaya Suresh. 2022. Online algorithms for matching platforms with multi-channel traffic. *Proceedings of the 23rd ACM Conference on Economics and Computation*. 986-987.
- Manshadi, Vahideh H, Shayan Oveis Gharan, Amin Saberi. 2012. Online stochastic matching: Online actions based on offline statistics. *Mathematics of Operations Research*, 37 (4), 559-573.
- Niazadeh, Rad, Negin Golrezaei, Joshua Wang, Fransisca Susan, Ashwinkumar Badanidiyuru. 2023. Online learning via offline greedy: Applications in market design and optimization. *Management Science*, 69 (7), 3797-3817.
- Ong, Qiyan, Jianying Qiu. 2023. Paying for randomization and indecisiveness. *Journal of Risk and Uncertainty*, 67 (1), 45-72.
- Rosário, Albérico Travassos, Joana Carmo Dias. 2023. How has data-driven marketing evolved: Challenges and opportunities with emerging technologies. *International Journal of Information Management Data Insights*, 3 (2), 100203.
- Rutten, Daan, Debankur Mukherjee. 2024. A new approach to capacity scaling augmented with unreliable machine learning predictions. *Mathematics of Operations Research*, 49 (1), 476-508.
- Shahmanzari, Masoud, Deniz Aksen. 2021. A multi-start granular skewed variable neighborhood tabu search for the roaming salesman problem. *Applied Soft Computing*, 102 107024.
- Shahmanzari, Masoud, Deniz Aksen, Saïd Salhi. 2020. Formulation and a two-phase matheuristic for the roaming salesman problem: Application to election logistics. *European Journal of Operational Research*, 280 (2), 656-670.
- Shahmanzari, Masoud, Renata Mansini. 2024. A learning-based granular variable neighborhood search for a multi-period election logistics problem with time-dependent profits. *European Journal of Operational Research*, 319 (1), 135-152.
- Shaw, Daron R. 1999. The effect of tv ads and candidate appearances on statewide presidential votes, 1988–96. *American Political Science Review*, 93 (2), 345-361.
- Shaw, Daron R. 2008. *The race to 270: The electoral college and the campaign strategies of 2000 and 2004*. University of Chicago Press.
- Shiri, Davood, Vahid Akbari, Ali Hassanzadeh. 2024. The capacitated team orienteering problem: An online optimization framework with predictions of unknown accuracy. *Transportation Research Part B: Methodological*, 185 102984.
- Tlili, Takwa, Saoussen Krichen. 2021. A simulated annealing-based recommender system for solving the tourist trip design problem. *Expert Systems with Applications*, 186 115723.
- Wang, Hao, Zhenzhen Yan, Xiaohui Bei. 2022. A nonasymptotic analysis for re-solving heuristic in online matching. *Production and Operations Management*, 31 (8), 3096-3124.