

Cesa-Bianchi, Nicolò; Colomboni, Roberto; Kasy, Maximilian

Working Paper

Adaptive Maximization of Social Welfare

IZA Discussion Papers, No. 17186

Provided in Cooperation with:

IZA – Institute of Labor Economics

Suggested Citation: Cesa-Bianchi, Nicolò; Colomboni, Roberto; Kasy, Maximilian (2024) : Adaptive Maximization of Social Welfare, IZA Discussion Papers, No. 17186, Institute of Labor Economics (IZA), Bonn

This Version is available at:

<https://hdl.handle.net/10419/302703>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

DISCUSSION PAPER SERIES

IZA DP No. 17186

Adaptive Maximization of Social Welfare

Nicolò Cesa-Bianchi
Roberto Colomboni
Maximilian Kasy

JULY 2024

DISCUSSION PAPER SERIES

IZA DP No. 17186

Adaptive Maximization of Social Welfare

Nicolò Cesa-Bianchi

Università degli Studi di Milano and Politecnico di Milano

Roberto Colomboni

Politecnico di Milano and Università degli Studi di Milano

Maximilian Kasy

University of Oxford and IZA

JULY 2024

Any opinions expressed in this paper are those of the author(s) and not those of IZA. Research published in this series may include views on policy, but IZA takes no institutional policy positions. The IZA research network is committed to the IZA Guiding Principles of Research Integrity.

The IZA Institute of Labor Economics is an independent economic research institute that conducts research in labor economics and offers evidence-based policy advice on labor market issues. Supported by the Deutsche Post Foundation, IZA runs the world's largest network of economists, whose research aims to provide answers to the global labor market challenges of our time. Our key objective is to build bridges between academic research, policymakers and society.

IZA Discussion Papers often represent preliminary work and are circulated to encourage discussion. Citation of such a paper should account for its provisional character. A revised version may be available directly from the author.

ISSN: 2365-9793

IZA – Institute of Labor Economics

Schaumburg-Lippe-Straße 5–9
53113 Bonn, Germany

Phone: +49-228-3894-0
Email: publications@iza.org

www.iza.org

ABSTRACT

Adaptive Maximization of Social Welfare*

We consider the problem of repeatedly choosing policies to maximize social welfare. Welfare is a weighted sum of private utility and public revenue. Earlier outcomes inform later policies. Utility is not observed, but indirectly inferred. Response functions are learned through experimentation.

We derive a lower bound on regret, and a matching adversarial upper bound for a variant of the Exp3 algorithm. Cumulative regret grows at a rate of $T^{2/3}$. This implies that (i) welfare maximization is harder than the multi-armed bandit problem (with a rate of $T^{1/2}$ for finite policy sets), and (ii) our algorithm achieves the optimal rate. For the stochastic setting, if social welfare is concave, we can achieve a rate of $T^{1/2}$ (for continuous policy sets), using a dyadic search algorithm.

We analyze an extension to nonlinear income taxation, and sketch an extension to commodity taxation. We compare our setting to monopoly pricing (which is easier), and price setting for bilateral trade (which is harder).

JEL Classification: C9, H21, C73

Keywords: optimal taxation, multi-armed bandits, experimental design

Corresponding author:

Maximilian Kasy
Department of Economics
University of Oxford Wellington Squar
Oxford
OX1 2JD
United Kingdom
E-mail: maximilian.kasy@economics.ox.ac.uk

* NCB was supported by the MUR PRIN grant 2022EKNE5K (Learning in Markets and Society) and by the FAIR (Future Artificial Intelligence Research) project, funded by the NextGenerationEU program within the PNRR-PE-AI scheme. Maximilian Kasy was supported by the Alfred P. Sloan Foundation, under the grant "Social foundations for statistics and machine learning."

1 Introduction

Consider a policymaker who aims to maximize social welfare, defined as a weighted sum of utility across individuals. The policymaker can choose a policy parameter such as a sales tax rate, an unemployment benefit level, a health-insurance copay rate, etc. The policymaker does *not* directly observe the welfare resulting from their policy choices. They do, however, observe behavioral outcomes such as consumption of the taxed good, labor market participation, or health care expenditures. They can revise their policy choices over time in light of observed outcomes. How should such a policymaker act? This is the question that we study. To address this question, we bring together insights from welfare economics (in particular optimal taxation, [Ramsey 1927](#); [Mirrlees 1971](#); [Bailey 1978](#); [Saez 2001](#); [Chetty 2009](#)) with insights from machine learning (in particular online learning and multi-armed bandits, see [Slivkins 2019](#); [Lattimore and Szepesvári 2020](#) for recent reviews, and [Thompson \(1933\)](#); [Lai and Robbins \(1985\)](#) for classic contributions).

In our baseline model, individuals arrive sequentially and make a single binary decision. In each period the policymaker chooses a tax rate that applies to this binary decision, and then observes the individual’s response. They do not observe the individual’s private utility. Social welfare is given by a weighted sum of private utility and public revenue. Later, we extend our model to nonlinear income taxation, where welfare weights vary as a function of individual earnings capacity, and sketch an extension to commodity taxation, where individual decisions involve a continuous consumption vector.

Our goal is to give guidance to the policymaker. We propose algorithms to maximize cumulative social welfare, and we provide (adversarial and stochastic) guarantees for the performance of these algorithms. In doing so, we also show that welfare maximization is a harder learning problem than reward maximization in the multi-armed bandit setting. Private utility in our baseline model is equal to consumer surplus, which is given by the integral of demand. In order to learn this integral, we need to learn demand for counterfactual, suboptimal tax rates. This drives the difficulty of the learning problem.

Why welfare, why adversarial guarantees? Our algorithms are designed to maximize social welfare, which is not directly observable, rather than maximizing outcomes that are directly observable. The definition of social welfare as an aggregation of individual utilities is at the heart of welfare economics in general, and of optimal tax theory

in particular. The distinction between utility and observable outcomes is important in practice. To illustrate, consider the example of a policymaker who chooses the level of unemployment benefits, where the observable outcome is employment. The policymaker could use an algorithm that adaptively maximizes employment. The problem with this approach is that employment might be maximized by making the unemployed as miserable as possible. This is not normatively appealing. Such an algorithm would *minimize* the utility of the unemployed, rather than maximizing social welfare. Similar examples can be given for many domains of public policy, including health, education, and criminal justice. In contrast to observable outcomes such as employment, welfare is improved by increasing the choice sets of those affected, not by reducing these choice sets.

Our theoretical analysis provides not only stochastic but also adversarial guarantees, which hold for arbitrary sequences of preference parameters. Adversarial guarantees for algorithms promise robustness against deviations from the assumption that heterogeneity is independently and identically distributed over time. Possible deviations from this assumption include autocorrelation, trends, heteroskedasticity, more general non-stationarity, and other concerns of time-series econometrics. In the employment example, such deviations might for instance be due to the business cycle. One might fear that adversarial robustness is achieved at the price of worsened performance for the i.i.d. setting, relative to less robust algorithms. That this is not the case follows from our theoretical characterizations.

Lower and upper bounds on regret Our main theorems provide lower and upper bounds on cumulative regret. Cumulative regret is defined as the difference in welfare between the *chosen* sequence of policies and the *best* possible constant policy. We consider both stochastic and adversarial regret. A lower bound on stochastic regret satisfies that, for any algorithm, there exists some stationary distribution of preference parameters for which the algorithm has to suffer at least a certain amount of regret. An upper bound on adversarial regret has to hold for a given algorithm and any sequence of preference parameters.

For a given algorithm, stochastic regret, averaged over i.i.d. sequences of preference parameters, is always less or equal than adversarial regret, for the worst-case sequence. A lower bound on stochastic regret (for any algorithm) therefore implies a corresponding lower bound on adversarial regret, and an upper bound on adversarial regret (for a given algorithm) immediately implies an upper bound on stochastic regret. When an

adversarial upper bound coincides with a stochastic lower bound, in terms of rates of regret, it follows that the proposed algorithm is rate efficient, for both stochastic and adversarial regret. It follows, furthermore, that the bounds are sharp.

A lower bound on stochastic regret We first prove a stochastic (and thus also adversarial) lower bound on regret, for any possible algorithm in the welfare maximization problem. Our proof of this bound constructs a family of possible distributions for preferences. This family is such that there are two candidate policies which are potentially optimal. The difference in welfare between these two policies depends on the integral of demand over intermediate policy values. In order to learn which of the two candidate policies is optimal, we need to learn behavioral responses for intermediate policies, which are strictly suboptimal. Because of the need to probe these suboptimal policies sufficiently often, we obtain a lower bound on regret which grows at a rate of $T^{2/3}$, even if we restrict our attention to settings with finite, known support for preference parameters and policies. This rate is worse than the worst-case rate for bandits of $T^{1/2}$.

A matching upper bound on adversarial regret for modified Exp3 We next propose an algorithm for the adaptive maximization of social welfare. Our algorithm is a modification of the Exp3 algorithm (Auer et al., 2002b). Exp3 is based on an unbiased estimate of cumulative welfare for each policy. The probability of choosing a given policy is proportional to the exponential of this estimate of cumulative welfare, times some rate parameter. Relative to Exp3, we require two modifications for our setting. First, we need to discretize the continuous policy space. Second, and more interestingly, we need additional exploration of counterfactual policies, including some policies that are clearly sub-optimal, in order to learn welfare for the policies which are contenders for the optimum. This need for additional exploration again arises because of the dependence of welfare on the integral of demand over counterfactual policy choices. For our modified Exp3 algorithm, we prove an adversarial (and thus also stochastic) upper bound on regret. We show that, for an appropriate choice of tuning parameters, worst case cumulative regret over all possible sequences of preference parameters grows at a rate of $T^{2/3}$, up to a logarithmic term. The algorithm thus achieves the best possible rate. Since the rates for our stochastic lower and adversarial upper bound coincide, up to a logarithmic term, we have a complete characterization of learning rates for the welfare maximization problem.

Improved stochastic bounds for concave social welfare The proof of our lower bound on regret is based on the construction of a distribution of preferences which delivers a non-concave social welfare function. If we restrict attention to the stochastic setting, where preferences are i.i.d. over time, and if we assume that social welfare is concave, then we can improve upon this bound on regret. We prove a lower bound on stochastic regret, under the assumption of concavity, which grows at the rate of $T^{1/2}$. We then propose a dyadic search algorithm which achieves this rate, up to logarithmic terms. This dyadic search algorithm maintains an “active interval,” containing the optimal policy with high probability, which is narrowed down over time. Only policies within the active interval are sampled.

Extensions to non-linear income taxation and to commodity taxation Our discussion up to this point focuses on a minimal, stylized case of an optimal tax problem, where individual actions are binary, and the policy imposes a tax on this binary action. Our arguments generalize, however, to more complicated and practically relevant settings. This includes optimal nonlinear income taxation, as in [Mirrlees \(1971\)](#) and [Saez \(2001\)](#), and commodity taxation for a bundle of goods, as in [Ramsey \(1927\)](#). For nonlinear income taxation, different tax rates apply at different income levels, and welfare weights depend on individual earnings capacity. In [Section 5](#), we discuss an extension of our tempered Exp3 algorithm to nonlinear income taxation, and characterize its regret. For commodity taxation, different tax rates apply to different goods, and consumption decisions are continuous vectors. In [Section 6](#) we sketch an extension of our algorithm to commodity taxation, but leave its characterization for future research.

Roadmap The rest of this paper proceeds as follows. We conclude this introduction with a discussion of some related work and relevant references. [Section 2](#) introduces our setup, formally defines the adversarial and stochastic settings, and compares our setup to related learning problems. [Section 3](#) provides lower and upper bounds on regret in the adversarial and stochastic settings. [Section 4](#) restricts attention to the stochastic setting with concave social welfare, and provides improved regret bounds for this setting. [Section 5](#) discusses an extension of our baseline model to non-linear income taxation. [Section 6](#) sketches another extension of our baseline model to commodity taxation. [Section 7](#) concludes, and discusses some possible applications of our algorithm, as well as an alternative Bayesian approach to adaptive welfare maximization. The proofs of [Theorem 1](#) and [Theorem 2](#) can be found in [Appendix A](#). The proofs of our remaining

theorems and proofs of technical lemmas are discussed in the Online Appendix.

1.1 Background and literature

To put our work in context, it is useful to contrast our framework with the standard approach in public finance and optimal tax theory, and with the frameworks considered in machine learning and the multi-armed bandit literature.

Optimal taxation Optimal tax theory, and optimal policy theory more generally, is concerned with the maximization of social welfare, where social welfare is understood as a (weighted) sum of subjective utility across individuals (Ramsey, 1927; Mirrlees, 1971; Baily, 1978; Saez, 2001; Chetty, 2009). A key tradeoff in such models is between, first, *redistribution* to those with higher welfare weights, and second, the efficiency cost of behavioral responses to tax increases. Such *behavioral responses* might reduce the tax base.

Optimal tax problems are defined by normative parameters (such as welfare weights for different individuals), as well as empirical parameters (such as the elasticity of the tax base with respect to tax rates). The typical approach in public finance uses historical or experimental variation to estimate the relevant empirical parameters (causal effects, elasticities). These estimated parameters are then plugged into formulas for optimal policy choice, which are derived from theoretical models. The implied optimal policies are finally implemented, without further experimental variation.

Multi-armed bandits The standard approach of public finance, which separates elasticity estimation from policy choice, contrasts with the adaptive approach that characterizes decision-making in many branches of AI, including online learning, multi-armed bandits, and reinforcement learning. Multi-armed bandit algorithms, in particular, trade off *exploration* and *exploitation* over time (Bubeck and Cesa-Bianchi, 2012; Slivkins, 2019; Lattimore and Szepesvári, 2020). Exploration here refers to the acquisition of information for better future policy decisions, while exploitation refers to the use of currently available information for optimal policy decisions at the present moment. The goal of bandit algorithms is to maximize a stream of rewards, which requires an optimal balance between exploration and exploitation. Bandit algorithms for the stochastic setting are characterized by optimism in the face of uncertainty: Policies with uncertain payoff should be tried until their expected payoff is clearly suboptimal.

Bandit algorithms (and similarly, adaptive experimental designs for informing policy choice, as in Russo [2020]; Kasy and Sautmann [2021]) are not directly applicable to social welfare maximization problems, such as those of optimal tax theory. The reason is that bandit algorithms maximize a stream of *observed* rewards. By contrast, social welfare as conceived in welfare economics is based on *unobserved* subjective utility.

Adversarial decision-making Adversarial models for sequential decision-making find their roots in repeated game theory (Hannan, 1957), while related settings were independently studied in information theory (Cover, 1965) and computer science (Vovk, 1990; Littlestone and Warmuth, 1994; Cesa-Bianchi et al., 1997). Regret minimization, also in a bandit setting, was investigated as a tool to prove convergence of uncoupled dynamics to equilibria in N -person games (Hart and Mas-Colell, 2000, 2001) – the exponential weighting scheme used by Exp3 is also known as *smooth fictitious play* in the game-theoretic literature (Fudenberg and Levine, 1995). Recent works (Seldin and Slivkins, 2014; Zimmert and Seldin, 2021) show that simple variants of Exp3 simultaneously achieve essentially optimal regret bounds in adversarial, stochastic, and contaminated settings, without prior knowledge of the actual regime. This suggests that algorithms designed for adversarial environments can behave well in more benign settings, whereas the opposite is provably not true.

Bandit approaches for economic problems Bandit-type approaches have been applied to a number of other economic and financial scenarios in the literature where rewards *are* observable. These include monopoly pricing (Kleinberg and Leighton, 2003) (see also the survey den Boer [2015]), second-price auctions (Cesa-Bianchi et al., 2015; Weed et al., 2016; Cesa-Bianchi et al., 2017), first-price auctions (Han et al., 2020b,a; Achddou et al., 2021; Cesa-Bianchi et al., 2024b)—see also (Kolumbus and Nisan, 2022; Feng et al., 2018, 2021), and combinatorial auctions (Daskalakis and Syrgkanis, 2022). Bandit-type approaches have also been applied to some settings where rewards are not directly observable, including bilateral trading (Cesa-Bianchi et al., 2021, 2023, 2024a), and the newsvendor problem (Lugosi et al., 2023).

Bandit algorithms are widely used in online advertising and recommendation. Online learning methods are successfully used for tuning the bids made by autobidders (a service provided by advertising platforms) (Lucier et al., 2024). While these algorithms are analyzed in adversarial environments, the extent to which they are deployed in commercial products remains unclear.

2 Setup

At each time $i = 1, 2, \dots, T$, one individual arrives who is characterized by an unknown willingness to pay $v_i \in [0, 1]$. This individual is exposed to a tax rate x_i , and makes a binary decision $y_i = \mathbf{1}(x_i \leq v_i)$. The implied public revenue is $x_i \cdot y_i$. The implied private welfare is $\max(v_i - x_i, 0)$. We define social welfare as a weighted sum of public revenue and private welfare, with a weight $\lambda \in (0, 1)$ for the latter. Social welfare for time period i is therefore given by

$$U_i(x_i) = \underbrace{x_i \cdot \mathbf{1}(x_i \leq v_i)}_{\text{Public revenue}} + \lambda \cdot \underbrace{\max(v_i - x_i, 0)}_{\text{Private welfare}}. \quad (1)$$

After period i , we observe y_i and the tax rate x_i , but nothing else. In particular, we do *not* observe welfare $U_i(x_i)$.

We can rewrite social welfare $U_i(x)$ as follows. Denote $G_i(x) = \mathbf{1}(v_i \geq x)$, so that $y_i = G_i(x_i)$. This is the individual demand function. Then private welfare can be written as $\max(v_i - x, 0) = \int_x^1 G_i(x') dx'$. That is, private welfare is given by integrated demand.¹ This representation of private welfare implies

$$U_i(x) = \underbrace{x \cdot G_i(x)}_{\text{Public revenue}} + \lambda \cdot \underbrace{\int_x^1 G_i(x') dx'}_{\text{Private welfare}}. \quad (2)$$

We consider algorithms for the choice of x_i which might depend on the observable history $(x_j, y_j)_{j=1}^{i-1}$, as well as possibly a randomization device.

Notation For the *adversarial* setting, we will consider cumulative demand and welfare, denoted by blackboard bold letters, summing across $j = 1, \dots, i$. In particular,

$$\mathbb{G}_i(x) = \sum_{j \leq i} G_j(x), \quad \mathbb{U}_i(x) = \sum_{j \leq i} U_j(x), \quad \mathbb{U}_i = \sum_{j \leq i} U_j(x_j).$$

$\mathbb{G}_i(x)$ and $\mathbb{U}_i(x)$ are cumulative demand and welfare for a counterfactual, fixed policy x . $\mathbb{U}_i$, without an argument, is the cumulative welfare for the policies x_j actually chosen.

For the *stochastic* setting, we will analogously consider expected demand and expected welfare, denoted by boldface letters. The expectation is taken across some

¹This reflects the absence of income effects in our model, which implies that private utility, consumer surplus, compensating variation, and equivalent variation all coincide.

stationary distribution μ of v_i , where v_i is statistically independent of x_i , and of v_j for $j \neq i$. In particular,

$$\mathbf{G}(x) = E[G_i(x)], \quad \mathbf{U}(x) = E[U_i(x)].$$

2.1 Regret

The adversarial case Following the literature, we consider regret for both the adversarial and the stochastic setting. In the adversarial setting, we allow for arbitrary sequences of willingness to pay, $\{v_i\}_{i=1}^T$. We compare the expected performance of any given algorithm for choosing $\{x_i\}_{i=1}^T$ to the performance of the best possible constant policy x . This comparison yields cumulative expected regret, which is given by

$$\mathcal{R}_T(\{v_i\}_{i=1}^T) = \sup_x E \left[\mathbb{U}_T(x) - \mathbb{U}_T \middle| \{v_i\}_{i=1}^T \right]. \quad (3)$$

The expectation in this expression is taken over any possible randomness in the tax rates x_i chosen by the algorithm; there is no other source of randomness.

The stochastic case We also consider the stochastic setting. In this setting, we add structure by assuming that the v_i are i.i.d. draws from some distribution μ on $[0, 1]$, with implied demand function $\mathbf{G}(x) = P(v_i \geq x)$. This demand function is identified by the regression

$$\mathbf{G}(x) = E[y_i | x_i = x].$$

The expectation in this expression is taken over the distribution of v_i , which is presumed to be statistically independent of the tax rate x_i . Expected welfare for this distribution of v_i is given by

$$\mathbf{U}(x) = x \cdot \mathbf{G}(x) + \lambda \int_x^1 \mathbf{G}(x') dx'.$$

Cumulative expected regret in the stochastic case equals

$$\mathcal{R}_T(\mathbf{G}) = \sup_x E [\mathbb{U}_T(x) - \mathbb{U}_T] = T \cdot \sup_x \mathbf{U}(x) - E \left[\sum_{i \leq T} U(x_i) \right]. \quad (4)$$

The expectation in this expression is taken over both any possible randomness in the tax rates x_i , and the i.i.d. draws v_i .

2.2 Comparison to related learning problems

Before proceeding with our analysis of regret, we take a step back, and compare our learning problem to two related problems that have received some attention in the literature. The first of these is the adaptive **monopoly pricing** problem; see for instance [Kleinberg and Leighton \(2003\)](#). This problem is equivalent to our setting when we set $\lambda = 0$, interpret x as a price, and U_i^{MP} as monopolist profits (neglecting production costs):

$$U_i^{\text{MP}}(x) = x_i \cdot \mathbf{1}(x_i \leq v_i) = \underbrace{x \cdot G_i(x)}_{\text{Monopolist revenue}}. \quad (5)$$

As in our adaptive taxation setting, the feedback received at the end of period i is

$$y_i = G_i(x_i) = \mathbf{1}(x_i \leq v_i).$$

Another related problem is price setting for **bilateral trade**, see for instance [Cesa-Bianchi et al. \(2024a\)](#). In this problem, welfare $U_i^{\text{BT}}(x)$ is given by the sum of seller and buyer welfare. Trade happens if and only if both sides agree to transact at the proposed price. Buyer willingness to pay is given by v_i^b , while the seller is willing to trade at prices above v_i^s .

$$\begin{aligned} U_i^{\text{BT}}(x) &= \mathbf{1}(v_i^b \geq x) \cdot \max(x - v_i^s, 0) + \mathbf{1}(v_i^s \leq x) \cdot \max(v_i^b - x, 0) \\ &= \underbrace{G_i^b(x) \cdot \int_0^x G_i^s(x') dx'}_{\text{Seller welfare}} + \underbrace{G_i^s(x) \cdot \int_x^1 G_i^b(x') dx'}_{\text{Buyer welfare}}. \end{aligned} \quad (6)$$

Feedback in this case is a little richer: We observe both whether the buyer b would have accepted the posted price, and whether the seller would have accepted this price,

$$y_i^b = G_i^b(x_i) = \mathbf{1}(x_i \leq v_i^b) \quad \text{and} \quad y_i^s = G_i^s(x_i) = \mathbf{1}(x_i \geq v_i^s).$$

Lipschitzness and information requirements The difficulty of the learning problem in each of these models critically depends on (i) the Lipschitz properties of the welfare function, and (ii) the information required to evaluate welfare at a point.

We say that a generic welfare function $W : [0, 1] \rightarrow \mathbb{R}$ is one-sided Lipschitz if $W(x + \varepsilon) \leq W(x) + \varepsilon$ for all $0 \leq x \leq 1$ and all $0 \leq \varepsilon \leq 1 - x$. One-sided Lipschitzness

Table 1: Regret rates for different learning problems

Model	Policy space		Objective function	
	Discrete	Continuous	Pointwise	One-sided Lipschitz
Monopoly price setting	$T^{1/2}$	$T^{2/3}$	Yes	Yes
Optimal taxation	$T^{2/3}$	$T^{2/3}$	No	Yes
Bilateral trade	$T^{2/3}$	T	No	No

Notes: This table shows the efficient rates of regret for different learning problems. Rates are up to logarithmic terms, and apply to both the stochastic and the adversarial setting. Regret rates are shown for the discrete case, where the space of policies x is restricted to a finite set, and the continuous case, where x can take any value in $[0, 1]$. The columns on the right describe the properties of the objective function in each problem, which drive the differences in regret rates.

Rates for the optimal taxation case are proven in this paper. Rates for the continuous monopoly price setting case are from [Kleinberg and Leighton \(2003\)](#); the discrete case reduces to a standard bandit problem. Rates for the continuous bilateral trade case are from [Cesa-Bianchi et al. \(2024a\)](#); the discrete case can be deduced by adapting the arguments in the same paper (for the stochastic i.i.d. case with independent sellers' and buyers' valuations), or by adapting the techniques in [Cesa-Bianchi et al. \(2023\)](#) (for the adversarial case, allowing the learner to use weakly budget balanced mechanisms).

allows us to bound the approximation error of a learning algorithm operating on a finite subset of the set of policies. One-sided Lipschitzness is an intrinsic property of both the monopoly pricing and the optimal taxation problem; it is not an assumption that is additionally imposed. To see this for monopoly pricing, note that, for $\epsilon \geq 0$, $U_i^{\text{MP}}(x + \epsilon) = (x + \epsilon) \cdot \mathbf{1}(x + \epsilon \leq v_i) \leq x \cdot \mathbf{1}(x \leq v_i) + \epsilon = U_i^{\text{MP}}(x) + \epsilon$. For social welfare, $U_i(x) = (x_i + \epsilon) \cdot \mathbf{1}(x_i + \epsilon \leq v_i) + \lambda \cdot \max(v_i - x_i - \epsilon, 0) \leq x \cdot \mathbf{1}(x \leq v_i) + \epsilon + \lambda \cdot \max(v_i - x_i, 0) = U_i(x) + \epsilon$.

We say that learning $W(\cdot)$ requires only pointwise information if $W(x)$ is a function of $G(x)$, and does not depend on $G(\cdot)$ otherwise. Pointwise information allows us to avoid exploring policies that are clearly suboptimal, when we aim to learn the optimal policy.

[Table 1](#) summarizes the Lipschitz properties and information requirements in each of the three models; the following justifies the claims made in [Table 1](#):

1. For **monopoly pricing**, welfare $U_i^{\text{MP}}(x)$ is one-sided Lipschitz and only depends on $G_i(x)$ pointwise.
2. For **optimal taxation**, welfare $U_i(x)$ is one-sided Lipschitz and depends on both $G_i(x)$ at the given x (pointwise), and on an integral of $G_i(x')$ for a range of values of x' (non-pointwise).

3. For **bilateral trade**, welfare $U_i^{\text{BT}}(x)$ is not one-sided Lipschitz and depends on both $G_i^b(x)$ and $G_i^s(x)$ (pointwise), as well as the integrals of $G_i^b(x')$ and $G_i^s(x')$ (non-pointwise).

These properties suggest a ranking in terms of the difficulty of the corresponding learning problems, and in particular in terms of the rates of divergence of cumulative regret: The information requirements of optimal taxation are stronger than those of monopoly pricing, but its continuity properties are more favorable than those of bilateral trade. This intuition is correct, as shown by [Table 1](#). The rates for monopoly pricing and for bilateral trade are known (or can be easily adapted) from the literature. In this paper we prove corresponding rates for optimal taxation.

In comparing optimal taxation and monopoly pricing to conventional multi-armed bandits, it is worth emphasizing that there are two distinct reasons for the slower rate of convergence. First, the continuous support of x , as opposed to a finite number of arms, which is shared by optimal taxation and monopoly pricing. Second, the requirement of additional exploration of sub-optimal policies for the optimal tax problem. As shown in [Table 1](#), the continuous support alone is enough to slow down convergence, with no extra penalty for the additional exploration requirement, in terms of rates. If, however, we restrict our attention to a discrete set of feasible policies x , then monopoly pricing reduces to a multi-armed bandit problem, with a minimax regret rate of $T^{1/2}$. The optimal tax problem, by contrast, still has a rate of $T^{2/3}$, even if we restrict our attention to the case of finite known support for v and x , as shown by the proof of [Theorem 1](#) below.

Hannan consistency The cumulative regret of any non-adaptive algorithm necessarily grows at a rate of T . This includes, in particular, randomized experiments where the policy is chosen uniformly at random, from a fixed policy set, in every period. Algorithms for which adversarial regret (and thus also stochastic regret) grows at a rate less than T , so that per-period regret goes to 0 as T increases, are known as *Hannan consistent*. Non-adaptive algorithms are not Hannan consistent. [Table 1](#) implies that Hannan consistent algorithms exist in all settings considered, with the exception of Bilateral trade and continuous policy spaces.

3 Stochastic and adversarial regret bounds

We now turn to our main theoretical results, lower and upper bounds on stochastic and adversarial regret for the problem of social welfare maximization. We first prove a lower bound on stochastic regret, which applies to any algorithm, and which immediately implies a lower bound on adversarial regret. We then introduce the algorithm Tempered Exp3 for Social Welfare. We show that, for an appropriate choice of tuning parameters, this algorithm achieves the rates of the lower bound on regret, up to a logarithmic term. Formal proofs of these bounds can be found in [Appendix A](#).

3.1 Lower bound

Theorem 1 (Lower bound on regret). *Consider the setup of Section [2](#). There exists a constant $C > 0$ such that, for any randomized algorithm for the choice of $x_1, x_2, \dots$ and any time horizon $T \in \mathbb{N}$, the following holds.*

1. *There exists a distribution μ on $[0, 1]$ with associated demand function $\mathbf{G}$ for which the stochastic cumulative expected regret $\mathcal{R}_T(\mathbf{G})$ is at least $C \cdot T^{2/3}$.*
2. *There exists a sequence $(v_1, \dots, v_T)$ for which the adversarial cumulative expected regret $\mathcal{R}_T(\{v_i\}_{i=1}^T)$ is at least $C \cdot T^{2/3}$.*

The proof of Theorem [1](#) can be found in [Appendix A](#). The adversarial lower bound follows immediately from the stochastic lower bound, since worst case regret (over possible sequences of v_i) is bounded below by average regret (over i.i.d. draws of v_i), for any distribution of v_i .

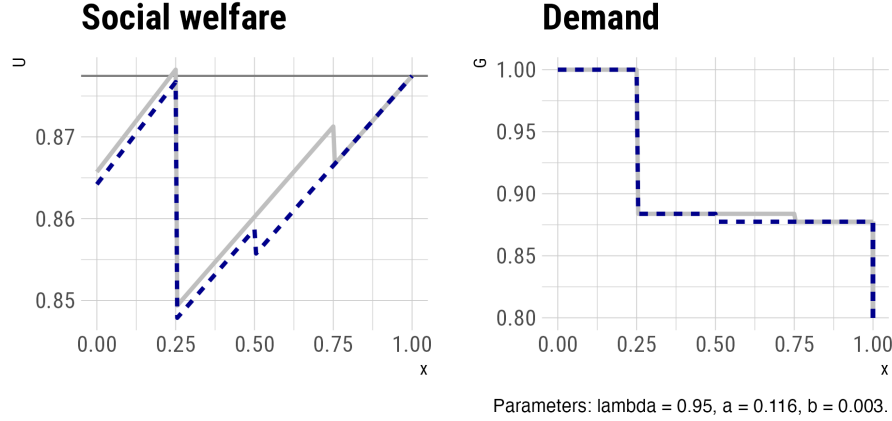
Sketch of proof To prove the stochastic lower bound we construct a family of distributions $\{\mu^\epsilon\}_{\epsilon \in [-1, 1]}$ for v_i , indexed by a parameter $\epsilon \in [-1, 1]$. The distributions in this family have four points of support, $(1/4, 1/2, 3/4, 1)$. The probability of these points is given by

$$(a, (1 + \epsilon)b, (1 - \epsilon)b, 1 - a - 2b).$$

The values of a and b are chosen such that (i) the two middle points $1/2, 3/4$ are far from optimal, for any value of ϵ , and (ii) learning which of the two end points $(1/4, 1)$ is optimal requires sampling from the middle.[2](#) For each $\epsilon \in [-1, 1]$, denote the demand

²Specifically, $a := \frac{(1-\lambda) \cdot (136-99\lambda)}{2 \cdot (4-3\lambda) \cdot (24-17\lambda)}$, and $b := \frac{1-\lambda}{2 \cdot (24-17\lambda)}$. These two constants are strictly greater than zero, and satisfy $1 - a - 2 \cdot b > 0$.

Figure 1: Construction for proving the lower bound on regret



Notes: This figure illustrates our construction for proving the lower bound on regret. The relative social welfare of policies 1 and .25 depends on the sign of ϵ . The bright line corresponds to $\epsilon = -1$, the dark, dashed line to $\epsilon = 1$. In order to distinguish between these two, we must learn demand in the intermediate interval $[.5, .75]$.

function associated to μ^ϵ by G^ϵ , and the expected social welfare associated to G^ϵ by U^ϵ . Property (ii) holds because of the integral term $\int_{1/4}^1 G^\epsilon(x') dx'$, which shows up in $U^\epsilon(1) - U^\epsilon(1/4)$. This construction is illustrated in [Figure 1](#). This figure shows plots of G^ϵ and of U^ϵ for $\lambda = .95$ and $\epsilon \in \{\pm 1\}$.

The difference in welfare $U^\epsilon(1) - U^\epsilon(1/4)$ of the two candidates optimal policies $1/4$ and 1 depends on the sign of ϵ . In order not to suffer expected regret that grows as $|\epsilon| \cdot T$, any learning algorithm needs to sample policies from points that are informative about the sign of ϵ . The only points that are informative are those in the region $(1/2, 3/4]$, where welfare is bounded away from optimal welfare.

More specifically, the learning algorithm has to sample on the order of $|\epsilon|^{-2}$ times from the region $(1/2, 3/4]$, to be able to detect the sign of ϵ , incurring regret on the order of $|\epsilon|^{-2}$ in the process. Any learning algorithm therefore incurs regret on the order of $\min(|\epsilon|^{-2}, |\epsilon| \cdot T)$, which, for $\epsilon \propto T^{-1/3}$, leads to the conclusion.

3.2 An algorithm that achieves the lower bound

We next introduce Algorithm [1](#), which allows us to essentially achieve the lower bound on regret, in terms of rates.

Algorithm 1 Tempered Exp3 for Social Welfare

Require: Tuning parameters K , γ and η .

- 1: Calculate evenly spaced grid-points $\tilde{x}_k = (k - 1)/K$,
and initialize $\widehat{\mathbb{G}}_{1k} = 0$ and $\widehat{\mathbb{U}}_{1k} = 0$ for $k = 1, \dots, K + 1$.
- 2: **for** individual $i = 1, 2, \dots, T$ **do**
- 3: For all $k = 1, 2, \dots, K + 1$, set {Assignment probabilities}

$$p_{ik} = (1 - \gamma) \cdot \frac{\exp(\eta \cdot \widehat{\mathbb{U}}_{ik})}{\sum_{k'} \exp(\eta \cdot \widehat{\mathbb{U}}_{ik'})} + \frac{\gamma}{K + 1}. \quad (7)$$

- 4: Choose k_i at random according to the probability distribution $(p_{i,1}, \dots, p_{i,K+1})$.
Set $x_i = \tilde{x}_{k_i}$, and query y_i accordingly.
- 5: For all $k = 1, 2, \dots, K + 1$, set {Estimated demand}

$$\widehat{\mathbb{G}}_{i+1,k} = \widehat{\mathbb{G}}_{i,k} + y_i \cdot \frac{\mathbf{1}(k_i = k)}{p_{ik}}. \quad (8)$$

- 6: For all $k = 1, 2, \dots, K + 1$, set {Estimated welfare}

$$\widehat{\mathbb{U}}_{i+1,k} = \tilde{x}_k \cdot \widehat{\mathbb{G}}_{i+1,k} + \frac{\lambda}{K} \cdot \sum_{k' > k} \widehat{\mathbb{G}}_{i+1,k'}. \quad (9)$$

7: **end for**

Conventional Exp3 Algorithm [1](#) is a modification of the Exp3 algorithm. Conventional Exp3 ([Auer et al., 2002b](#)) is designed to maximize the standard bandit objective, $\sum_{i \leq T} y_i$. Exp3 maintains an unbiased running estimate of the cumulative payoff of each arm k , calculated using inverse probability weighting, $\widehat{\mathbb{G}}_{i,k} = \sum_{j < i} y_j \cdot \frac{\mathbf{1}(k_j = k)}{p_{jk}}$. In period i , arm k is chosen with probability $p_{ik} = (1 - \gamma) \cdot \frac{\exp(\eta \cdot \widehat{\mathbb{G}}_{ik})}{\sum_{k'} \exp(\eta \cdot \widehat{\mathbb{G}}_{ik'})} + \frac{\gamma}{K+1}$, where η and γ are tuning parameters. p_{ik} is thus increasing in the estimated average performance $\frac{\widehat{\mathbb{G}}_{i,k}}{i}$ of arm k in prior periods. Because $\widehat{\mathbb{G}}_{i,k}$ is *not* normalized by the number of time periods k , more weight is given to the best-performing arms over time, as estimation uncertainty for average performance decreases. In both these aspects, Exp3 is similar to the popular Upper Confidence Bound algorithm (UCB) for stochastic bandit problems ([Lai, 1987](#); [Agrawal, 1995](#); [Auer et al., 2002a](#)). In contrast to UCB, Exp3 is a randomized algorithm. Randomization is required for adversarial performance guarantees. This is analogous to the necessity of mixed strategies in Nash equilibrium for zero-sum games.

Modifications relative to conventional Exp3 Relative to this algorithm, we require three modifications. First, we discretize the continuous support $[0, 1]$ of x , restricting attention to the grid of policy values $\tilde{x}_k = (k - 1)/K$. Second, since welfare $U_i(x)$ is not directly observed for the chosen policy x , we need to estimate it indirectly. In particular, we first form an estimate $\hat{\mathbb{G}}_{ik}$ of cumulative demand for each of the policy values $\tilde{x}_k$, using inverse probability weighting. We then use this estimated demand, interpolated using a step-function, to form estimates of cumulative social welfare, $\hat{\mathbb{U}}_{ik} = \tilde{x}_k \cdot \hat{\mathbb{G}}_{ik} + \frac{\lambda}{K} \cdot \sum_{k' > k} \hat{\mathbb{G}}_{ik'}$. Third, we require additional exploration, relative to Exp3. Since social welfare depends on demand for counterfactual policy choices, we need to explore policies that are away from the optimum, in order to learn the relative welfare of approximately optimal policy choices. The mixing weight γ , which determines the share of policies sampled from the uniform distribution, needs to be larger relative to conventional Exp3, to ensure sufficient exploration away from the optimum.

Theorem 2 (Adversarial upper bound on regret of Tempered Exp3 for Social Welfare).

Consider the setup of Section 2, and Algorithm 1. Assume that $(K + 1)\eta < \gamma$.

Then for any sequence $(v_1, \dots, v_T)$ expected regret $\mathcal{R}_T(\{v_i\}_{i=1}^T)$ is bounded above by

$$\left(\gamma + \eta \cdot (e - 2)^{\frac{K+1}{K}} \cdot \left(\frac{2K+1}{6} + \frac{\lambda^2}{\gamma} \right) + \frac{\lambda}{K} \right) \cdot T + \frac{\log(K+1)}{\eta}. \quad (10)$$

Suppose additionally that $c_1, c_2, c_3 > 0$ are constants. Then, there exists a constant c_4 such that, if we set $\gamma = c_1 \cdot \left(\frac{\log(T)}{T} \right)^{1/3}$, $\eta = c_2 \cdot \gamma^2$, and $K = \lfloor c_3/\gamma \rfloor$, the expected regret $\mathcal{R}_T(\{v_i\}_{i=1}^T)$ is bounded above by

$$c_4 \cdot \log(T)^{1/3} T^{2/3}. \quad (11)$$

Corollary 1 (Stochastic upper bound on regret of Tempered Exp3 for Social Welfare).

Under the assumptions of Theorem 2, suppose additionally that v_i is drawn i.i.d. from some distribution with associated demand function $\mathbf{G}$. Then expected regret $\mathcal{R}_T(\mathbf{G})$ is bounded above by the same expressions as in Theorem 2.

The proof of Theorem 2 can again be found in Appendix A.

Tuning The statement of the theorem leaves the constants c_1, c_2, c_3 in the definition of the tuning parameters unspecified. Suppose we wish to choose the tuning parameters so as to optimize the upper bound obtained in Theorem 2. Ignoring the rounding of

K , an approximate solution to this problem is given by

$$\begin{aligned}\eta &= 1/a \cdot (\log(T)/T)^{2/3} \\ \gamma &= \lambda \sqrt{(e-2)/a} \cdot (\log(T)/T)^{1/3} \\ K &= \sqrt{3\lambda a/(e-2)} \cdot (T/\log(T))^{1/3}\end{aligned}$$

where

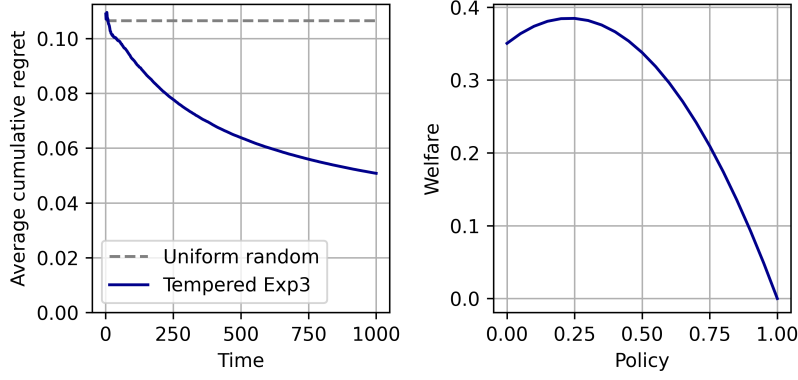
$$a = (9(e-2))^{1/3} (\sqrt{\lambda/3} + \lambda)^{2/3}.$$

This solution is obtained by taking the upper bound in Equation (10), approximating $(K+1)/K \approx 1$ and $(2K+1)/6 \approx K/3$, and solving the first order conditions with respect to the three tuning parameters. This approximation, and the tuning parameters specified above, yield an approximate upper bound on regret of $6 \cdot \log(T)^{1/3} T^{2/3}$.

Unknown time horizon Note that the proposed tuning depends crucially on knowledge of the time horizon T at which regret is to be evaluated. In order to extend our rate results to the case of unknown time horizons, we can use the so-called doubling trick; cf. Section 2.3 of [Cesa-Bianchi and Lugosi \(2006\)](#): Consider a sequence of epochs (intervals of time-periods) of exponentially increasing length, and re-run Algorithm 1 for each time-period separately, tuning the parameters over the current epoch length. This construction converts Algorithm 1 into an “anytime algorithm” which enjoys the same regret guarantees of Theorem 2, up to a multiplicative constant factor. Another more efficient strategy to achieve the same goal is to modify Algorithm 1, allowing the parameters η and γ to change at each iteration, and splitting each bin associated with the discretization parameter K whenever more precision is required.

The extra $\log(T)^{1/3}$ term There is a rate discrepancy between our upper and lower bounds on regret, corresponding to the $\log(T)^{1/3}$ term in our upper bound. We conjecture the existence of an alternative algorithm that can eliminate this extra logarithmic term, albeit at the cost of reduced computational efficiency and a less transparent theoretical analysis. Our conjecture is based on known results for the standard multi-armed bandit problem with K arms. The Exp3 algorithm achieves an upper bound of order $\sqrt{K \log(K) T}$ for this problem, which includes an extra logarithmic term compared to the known lower bound of order $\sqrt{KT}$. Exp3 is an instance of the Follow-The-Regularized-Leader (FTRL) algorithm with importance weighting and the

Figure 2: Tempered Exp3 for Social Welfare– Numerical example



Notes: This figure illustrates the performance of our algorithm for the stochastic case, where v_i is drawn uniformly at random from $[0, 1]$ for all i , the weight λ equals .7, and the tuning parameters are $K = 20$, $\eta = .025$, $\gamma = .1$. The left plot shows the cumulative average regret of our algorithm, averaged across 4000 simulations. The right plot shows expected social welfare $U(x)$ as a function of the policy x .

negative entropy as the regularizer. It is known that using the $\frac{1}{2}$ -Tsallis entropy as the regularizer in the FTRL algorithm with importance weighting results in regret guarantees of order $\sqrt{KT}$ for the bandit problem (Lattimore and Szepesvári, 2020). However, unlike Exp3, the FTRL with Tsallis entropy involves a more complex proof and requires solving an optimization problem in each period. Analogous statements might be true for our setting.

Numerical example For illustration, Figure 2 plots the cumulative average regret of Tempered Exp3 for Social Welfare for the case where v_i is sampled uniformly at random from $[0, 1]$ each time period. Initially, the performance of our algorithm is, by construction, equal to the performance of choosing a policy uniformly at random. Over time, however, the average regret of our algorithm drops by more than half, in this numerical example. Note that the rate at which cumulative regret declines in Figure 2 (for i.i.d. sampling from a fixed distribution) is unrelated to the regret rate of Theorem 2 (for the worst case sequence of v_i , for each time horizon T).

4 Stochastic regret bounds for concave social welfare

Theorem 1 in Section 3 provides a lower bound proportional to $T^{2/3}$ for adversarial and stochastic regret in social welfare maximization. The proof of this lower bound constructs a distribution for the v_i . This distribution is such that expected social welfare $U(x)$ is non-concave, as a function of x ; two global optima are separated by a region of lower welfare. In order to learn which of two candidates for the globally optimal policy is actually optimal, it is necessary to sample policies in between. These intermediate policies yield lower welfare, and sampling them contributes to cumulative regret. This construction is illustrated in Figure 1.

Given that the construction relies on non-concavity of expected social welfare, could we achieve lower regret if we knew that social welfare is actually concave? The answer turns out to be yes, for the stochastic setting (in the adversarial setting, cumulative welfare is necessarily non-concave). One reason is that concavity ensures that the function is unimodal. To estimate the difference in social welfare between two policies it therefore suffices to sample policies that lie in the interval between them. These in-between policies yield social welfare exceeding the minimum of the two boundary policies. A second reason is that concavity prevents unexpected spikes in social welfare. This property allows us to test carefully chosen triples of points for extended periods, to ensure that one of them is suboptimal, without incurring significant regret.

For the stochastic setting with concave social welfare, we present an algorithm that achieves a bound on regret of order $T^{1/2}$, up to logarithmic terms. Before describing our proposed algorithm, Dyadic Search for Social Welfare, let us formally state the improved regret bounds. The proofs of these lower and upper bounds can be found in Online Appendix B.

Theorem 3 (Lower bound on regret for the concave case). *Consider the setup of Section 2. There exists a constant $C > 0$ such that, for any randomized algorithm for the choice of $x_1, x_2, \dots$ and any time horizon $T \in \mathbb{N}$, the following holds:*

There exists a distribution μ on $[0, 1]$ with associated demand function $\mathbf{G}$ and concave social welfare function $\mathbf{U}$, for which the stochastic cumulative expected regret $\mathcal{R}_T(\mathbf{G})$ is at least $C \cdot T^{1/2}$.

Theorem 4 (Stochastic upper bound on regret of Dyadic Search for Social Welfare). *Consider the stochastic setup of Section 2 and time horizon $T \in \mathbb{N}$. If Algorithm 2 is*

Algorithm 2 Dyadic Search for Social Welfare

Require: A confidence parameter $\delta \in (0, 1)$.

```
1:  $I_1 = [0, 1]$ ,  $t_0 = 0$ ,  $k = 0$ 
2: for epochs  $\tau = 1, 2, \dots$  do
3:   Let  $c = (\sup I_\tau + \inf I_\tau)/2$ , and  $d = \sup I_\tau - \inf I_\tau$ . {Subinterval for sampling}
4:   if  $\tau$  is odd then
5:     Let  $l = c - \frac{1}{4}d$ ,  $r = c + \frac{1}{4}d$ .
6:   else
7:     Let  $l = c - \frac{1}{6}d$ ,  $r = c + \frac{1}{6}d$ .
8:   end if
9:   for  $t = t_{\tau-1} + 1, t_{\tau-1} + 2, \dots$  do
10:    Select  $w \in \operatorname{argmax}_{w' \in \{l, c, r, (l, c), (c, r)\}} \Gamma_{t-1}(w')$ , {Sampling}
        breaking ties following the order  $l, c, r, (l, c), (c, r)$ 
11:    if  $w \in \{l, c, r\}$  then
12:      Set  $x_t = w$ .
13:    else
14:      Set  $x_t = w_1 + (w_2 - w_1) \cdot \frac{k+1/2}{n_{t-1}(w_1, w_2)+1}$ , and  $k = (k+1) \bmod n_{t-1}(w_1, w_2) + 1$ .
15:    end if
16:    Calculate  $J_t(l, c)$ ,  $J_t(c, r)$ , and  $J_t(l, r)$ , as in Equations (15) and (16). {Inference}
17:    if  $\inf(J_t(l, c)) \geq 0$  or  $\inf(J_t(l, r)) \geq 0$  then
18:      let  $I_{\tau+1} = I_\tau \cap [l, 1]$  and  $t_\tau = t$  and break {Shrinking the active interval}
19:    else if  $\sup(J_t(c, r)) \leq 0$  or  $\sup(J_t(l, r)) \leq 0$  then
20:      let  $I_{\tau+1} = I_\tau \cap [0, r]$  and  $t_\tau = t$  and break
21:    end if
22:  end for
23: end for
```

run with confidence parameter $\delta = \frac{1}{T^{5/2}}$, and if the social welfare function $\mathbf{U}$ is concave, then, the expected regret $\mathcal{R}_T(\mathbf{G})$ is of order at most $T^{1/2}$, up to logarithmic terms.

Dyadic search Our algorithm is based on a modification of dyadic search, as discussed in (Bachoc et al., 2022a,b). At any point in time, this algorithm maintains an active interval I_τ , which contains the optimal policy with high probability. Only policies within this interval are sampled going forward. As evidence accumulates, this interval is trimmed down, by excluding policies that are sub-optimal with high probability.

The algorithm proceeds in epochs τ . At the start of each epoch, a sub-interval $[l, r] \subset I_\tau$ is formed, with mid-point $c = (l+r)/2$. The points l, c, r are in a dyadic grid, that is, they are of the form $k/2^m$. After sampling from $[l, r]$, we calculate confidence

intervals $J_t(l, c)$, $J_t(c, r)$, and $J_t(l, r)$ for the welfare differences $\Delta(l, c)$, $\Delta(c, r)$, and $\Delta(l, r)$, where $\Delta(x, x') = \mathbf{U}(x') - \mathbf{U}(x)$.

If the confidence interval $J_t(l, c)$ or $J_t(l, r)$ lies above 0, concavity implies that the optimal policy cannot lie to the left of l ; we can thus trim the active interval I_τ by dropping all points to the left of l . Symmetrically, if the confidence interval $J_t(c, r)$ or $J_t(l, r)$ lies below 0, we can trim I_τ by dropping all points to the right of r .

Confidence intervals for welfare differences This procedure requires the construction of confidence intervals for welfare differences of the form

$$\Delta(x, x') = \mathbf{U}(x') - \mathbf{U}(x) = x' \cdot \mathbf{G}(x') - x \cdot \mathbf{G}(x) - \lambda \int_x^{x'} \mathbf{G}(x'') dx''. \quad (12)$$

At time t , we estimate demand $\mathbf{G}(x)$, for policies x chosen in previous periods, as³

$$\hat{\mathbf{G}}_t(x) = \frac{1}{n_t(x)} \sum_{i \leq t} y_i \cdot \mathbf{1}(x_i = x), \quad n_t(x) = \sum_{i \leq t} \mathbf{1}(x_i = x).$$

We similarly estimate integrated demand $\int_x^{x'} \mathbf{G}(x'') dx''$ by $(x' - x)$ times the average of realized demand y_i for observations x_i in the open interval (x, x') . We have to be careful, however, to use a sample of x_i that is (approximately) uniformly distributed over this interval. This can be achieved for our dyadic search procedure, as specified in Algorithm 2, by truncating the time index used to estimate this average.⁴ Let

$$s(x, x', t) = \max \left\{ s \leq t : \log_2 \left(1 + \sum_{i \leq s} \mathbf{1}(x_i \in (x, x')) \right) \in \mathbb{N} \right\}.$$

We define

$$\hat{\mathbf{G}}_t(x, x') = \frac{1}{n_t(x, x') + 1} \sum_{i \leq s(x, x', t)} y_i \cdot \mathbf{1}(x_i \in (x, x')), \quad n_t(x, x') = \sum_{i \leq s(x, x', t)} \mathbf{1}(x_i \in (x, x')).$$

At each round, Algorithm 2 maintains estimates for welfare differences among three points l, c, r (for left, center and right, respectively). The estimate of the welfare dif-

³We use the convention $0/0 = 0$ and $a/0 = +\infty$ whenever $a > 0$. Furthermore, every summation over an empty set of indices is understood to have value 0.

⁴The sampling procedure in Algorithm 2 samples sequentially from the dyadic grid in the active interval, refining the grid in subsequent iterations. $s(x, x', t)$ provides a truncation of the time index such that one round of such dyadic sampling has been completed.

ference between $x' = c$ and $x = l$ (or between $x' = r$ and $x = c$) is given by

$$\widehat{\Delta}_t(x, x') = x' \cdot \widehat{\mathbf{G}}_t(x') - x \cdot \widehat{\mathbf{G}}_t(x) - \lambda \cdot (x' - x) \cdot \widehat{\mathbf{G}}_t(x, x'). \quad (13)$$

while the estimate of the welfare difference between r and l is given by

$$\widehat{\Delta}_t(l, r) = \widehat{\Delta}_t(l, c) + \widehat{\Delta}_t(c, r). \quad (14)$$

To construct confidence intervals for $\Delta(x, x')$, we also need to quantify the uncertainty of our demand estimates. We use the following interval half-lengths for confidence intervals for tax revenue at x , and for the private welfare difference between x' and x :

$$\Gamma_t(x) = x \cdot \sqrt{\frac{1}{2n_t(x)} \log\left(\frac{2}{\delta}\right)}, \quad \Gamma_t(x, x') = \lambda \cdot (x' - x) \cdot \left(\sqrt{\frac{1}{2(n_t(x, x') + 1)} \log\left(\frac{2}{\delta}\right)} + \frac{2}{n_t(x, x') + 1} \right).$$

Using the shorthand $a \pm b = [a - b, a + b]$, our confidence interval for $\Delta(x, x')$, where $x' = c$ and $x = l$ (or $x' = r$ and $x = c$) is given by

$$J_t(x, x') = \widehat{\Delta}_t(x, x') \pm (\Gamma_t(x') + \Gamma_t(x) + \Gamma_t(x, x')), \quad (15)$$

while our confidence interval for $\Delta(l, r)$ is given by

$$J_t(l, r) = \widehat{\Delta}_t(l, r) \pm (\Gamma_t(r) + \Gamma_t(l) + \Gamma_t(l, c) + \Gamma_t(c, r)). \quad (16)$$

With these preliminaries, we are now ready to state our algorithm, Dyadic Search for Social Welfare, in Algorithm [2](#).

Before concluding this section, we highlight two features of Algorithm [2](#). First, two of the three points l, c, r , and the corresponding estimates of demand, are kept from each epoch to the next. Second, estimation of the integral term is performed by querying points following a fixed and balanced design on the dyadic grid – instead of, for example, using a randomized Monte Carlo procedure which may lead to unbalanced exploration. This implies that the points queried to estimate the integral terms can be easily reused to obtain other integral estimates from each epoch to the next. These two features combined ensure that Algorithm [2](#) recycles information very efficiently to prune the active interval as quickly as possible, which leads to better regret.

5 Income taxation

We discuss two extensions of the baseline model of optimal taxation that we introduced in Section 2. These extensions incorporate features that are important in more realistic models of optimal taxation. For both of these extensions, we propose a properly modified version of Algorithm 1. The first extension, discussed in this section, is a variant of the Mirrlees model of optimal income taxation (Mirrlees, 1971; Saez, 2001, 2002). The second extension, discussed in Section 6 is a variant of the Ramsey model of commodity taxation (Ramsey, 1927).

Our model of income taxation generalizes our baseline model by allowing for heterogeneous wages w_i , welfare weights $\omega(w_i)$, extensive-margin labor supply responses determined by the cost of participation v_i , and non-linear income taxes $x_i = \mathbf{x}(w_i)$. Two simplifications are maintained in this model, relative to a more general model of income taxation. First, only extensive margin responses (participation decisions) by individuals are allowed; there are no intensive margin responses (hours adjustments). Second, as in the baseline model of Section 2, there are no income effects. In imposing these assumptions, our model mirrors the model of optimal income taxation discussed in Section II.2 of Saez (2002).

Setup At each time $i = 1, 2, \dots, T$, one individual arrives who is characterized by (i) a potential wage $w_i \in [0, 1]$, and (ii) an unknown cost of participation $v_i \in [0, 1]$. This individual makes a binary labor supply decision y_i . If they participate in the labor market ($y_i = 1$), they earn w_i , but pay a tax according to the tax rate $x_i = \mathbf{x}(w_i)$ on their earnings w_i . They furthermore incur a non-monetary cost of participation v_i .

Their optimal labor supply decision is therefore given by $y_i = \mathbf{1}(v_i \leq w_i \cdot (1 - x_i))$, and private welfare equals $\max(w_i \cdot (1 - x_i) - v_i, 0)$. The implied public revenue is equal to the tax on earnings $x_i \cdot w_i$ if $y_i = 1$, and 0 otherwise.

We define social welfare as a weighted sum of public revenue and private welfare, with a weight $\omega(w_i)$ for the latter. Typically, ω is a decreasing function of w , reflecting a preference for redistribution towards those with lower earnings potential, cf. Saez and Stantcheva (2016). Social welfare for time period i , as a function of the tax schedule $\mathbf{x}(\cdot)$, is therefore given by

$$U_i(\mathbf{x}(\cdot)) = \underbrace{\mathbf{x}(w_i) \cdot w_i \cdot \mathbf{1}(v_i \leq w_i \cdot (1 - \mathbf{x}(w_i)))}_{\text{Public revenue}} + \underbrace{\omega(w_i) \cdot \max(w_i \cdot (1 - \mathbf{x}(w_i)) - v_i, 0)}_{\text{Private welfare}}. \quad (17)$$

After period i , we observe y_i and the tax schedule $\mathbf{x}_i(\cdot)$. If $y_i = 1$, we also observe w_i . Nothing else is observed.⁵

Piecewise constant tax schedules We next construct a generalization of Algorithm 1 based on piecewise constant tax schedules, with tax rates changing at the grid-points $\mathcal{W}$, where $0 \in \mathcal{W} \subset [0, 1]$. Formally, define $\lfloor w \rfloor = \max\{w' \in \mathcal{W} : w' \leq w\}$, rounding the wage w down to the nearest grid-point in $\mathcal{W}$.⁶ Denote $H = |\mathcal{W}|$, and let

$$\mathcal{X}_{\mathcal{W}} = \{\mathbf{x}(\cdot) : \forall w \in [0, 1], \mathbf{x}(w) = \mathbf{x}(\lfloor w \rfloor)\}.$$

For $w \in \mathcal{W}$ and any $x \in [0, 1]$, denote

$$G_i(w, x) = w_i \cdot \mathbf{1}(v_i \leq w_i \cdot (1 - x)) \cdot \mathbf{1}(\lfloor w_i \rfloor = w),$$

so that $y_i \cdot w_i = G_i(w_i, \mathbf{x}_i(w_i))$. $G_i(w, x)$ is the individual labor supply function, in monetary units, interacted with an indicator for whether the wage w_i falls into the tax bracket starting at w . With this notation, we can rewrite

$$\max(w_i \cdot (1 - x) - v_i, 0) = \int_x^1 G_i(\lfloor w_i \rfloor, x') dx'.$$

For piecewise constant tax rates $\mathbf{x}(\cdot)$ we get

$$U_i(\mathbf{x}(\cdot)) = \sum_{w \in \mathcal{W}} \left[\mathbf{x}(w) \cdot G_i(w, \mathbf{x}(w)) + \omega(w_i) \cdot \int_{\mathbf{x}(w)}^1 G_i(w, x') dx' \right]. \quad (18)$$

Cumulative social welfare is given by $\mathbb{U}_i = \sum_{j \leq i} U_i(\mathbf{x}_i(\cdot))$, and we correspondingly define cumulative expected regret, in the adversarial setting, as

$$\mathcal{R}_T = \sup_{\mathbf{x}(\cdot) \in \mathcal{X}_{\mathcal{W}}} E \left[\mathbb{U}_T(\mathbf{x}(\cdot)) - \mathbb{U}_T \left| \{v_i\}_{i=1}^T, \{w_i\}_{i=1}^T \right. \right].$$

⁵It should be noted that in this model we take the transfer x_0 for individuals without other income as given. The effective tax owed by an employed individual equals $\mathbf{x}(w_i) \cdot w_i - x_0$. The “unconditional basic income” x_0 does not affect labor supply, given our assumption that there are no income effects, and it enters social welfare additively. It is therefore without loss of generality to omit x_0 from our model.

⁶Here we use slightly non-standard notation, where $\lfloor \cdot \rfloor$ denotes rounding down to the nearest grid-point, rather than the nearest integer.

The supremum here is taken over all tax schedules $\mathbf{x}(\cdot)$ that are piecewise constant between the gridpoints $w \in \mathcal{W}$.

Algorithm Algorithm 3 generalizes Algorithm 1 to this setting. As before, we form an unbiased estimate $\hat{G}_i$ of G_i using inverse probability weighting, map this estimate into a corresponding estimate $\hat{U}_i$ of U_i , based on Equation (18), and cumulate across time periods to obtain $\hat{U}_i$. Note that w_i is observed whenever $y_i = 1$. This implies that the estimate $\hat{G}_i$ is in fact a function of observables, and the same holds for $\hat{U}_i$.

Algorithm 3 keeps track of estimated demand and social welfare for each bin (“tax bracket”), as defined by the gridpoints $w \in \mathcal{W}$. The algorithm then constructs a distribution $p_i(x|w)$ over tax rates $x \in \mathcal{X}$ given w , using the tempered Exp3 distribution. The tax schedule $\mathbf{x}(\cdot)$ is sampled according to these (marginal) distributions of tax rates for each bracket. Though immaterial for the following theorem, we choose the perfectly correlated coupling, across brackets, of these marginal distributions, which is implemented using the random variable A_i in Algorithm 3.

Theorem 5 (Adversarial upper bound on regret of Tempered Exp3 for Optimal Income Taxation). *Consider the setup of Section 5, and Algorithm 3. Assume that $(K+1)\eta < \gamma$, and that $\omega(w) \leq 1$ for all w .*

Then for any sequence $(v_1, \dots, v_T)$ expected regret $\mathcal{R}_T(\{v_i\}_{i=1}^T)$ is bounded above by

$$\left(\gamma + \eta \cdot (e - 2) \frac{K+1}{K} \cdot \left(\frac{2K+1}{6} + \frac{1}{\gamma}\right) + \frac{1}{K}\right) \cdot T + \frac{H \log(K+1)}{\eta}. \quad (23)$$

Suppose additionally⁷ that $K = c_1 \cdot (T/H)^{1/3}$, $\gamma = c_2/(K+1)$, and $\eta = c_3/(K+1)^2$, for some constants c_1, c_2, c_3 . Then expected regret $\mathcal{R}_T(\{v_i\}_{i=1}^T)$ is bounded above by

$$c_4 \cdot H^{1/3} \cdot \log(T)^{1/3} T^{2/3}, \quad (24)$$

for some constant c_4 .

6 Commodity taxation

In this section, we generalize our baseline model of optimal taxation to a model of commodity taxation with multiple goods $j \in \{1, \dots, k\}$ and continuous demand functions

⁷for simplicity, we assume that in the following tuning K is an integer. If not, round K to the closest integer.

Algorithm 3 Tempered Exp3 for Optimal Income Taxation

Require: Tuning parameters K , γ and η , and set of gridpoints $\mathcal{W} \subset [0, 1]$.

- 1: Calculate evenly spaced grid-points $\mathcal{X} = \{0, \frac{1}{K}, \frac{2}{K}, \dots, 1\}$.
- 2: Initialize $\widehat{G}_1(w, x) = 0$ and $\widehat{U}_1(w, x) = 0$ for all $w \in \mathcal{X}$ and all $x \in \mathcal{X}$.
- 3: **for** individual $i = 1, 2, \dots, T$ **do**
- 4: For all $x, w \in \mathcal{X}$, set $\lfloor w \rfloor = \max\{w' \in \mathcal{W} : w' \leq w\}$, and
{Assignment probabilities}

$$p_i(x|w) = (1 - \gamma) \cdot \frac{\exp(\eta \cdot \widehat{U}_i(x, \lfloor w \rfloor))}{\sum_{x' \in \mathcal{X}} \exp(\eta \cdot \widehat{U}_i(x', \lfloor w \rfloor))} + \frac{\gamma}{K + 1}. \quad (19)$$

- 5: Draw $A_i \sim U[0, 1]$. For all $w \in [0, 1]$, set

$$\mathbf{x}_i(w) = \max \left\{ x \in \mathcal{X} : \sum_{x' \in \mathcal{X}, x' < x} p_i(x'|w) \leq A_i \right\}, \quad (20)$$

and query y_i accordingly.

- 6: For all $w \in \mathcal{W}$ and $x \in \mathcal{X}$, set {Estimated labor supply}

$$\widehat{G}_i(x, w) = y_i \cdot w_i \cdot \frac{\mathbf{1}(\lfloor w_i \rfloor = w, \mathbf{x}_i(w_i) = x)}{p_i(x|w)}. \quad (21)$$

- 7: For all $w \in \mathcal{W}$ and $x \in \mathcal{X}$, set {Estimated welfare}

$$\widehat{U}_{i+1}(x, w) = \widehat{U}_i(x, w) + x \cdot \widehat{G}_i(x, w) + \frac{\omega(w_i)}{K} \cdot \sum_{x' \in \mathcal{X}, x' > x} \widehat{G}_i(x', w). \quad (22)$$

8: **end for**

$y_i(x) \in [0, 1]^k$, where $x \in [0, 1]^k$ is a vector of tax rates. We again assume that there are no income effects. Our setup is a version of the classic Ramsey model (Ramsey, 1927). We propose a generalization of Tempered Exp3 for Social Welfare to this setting. In the following, we use $\langle x, y \rangle$ to denote the Euclidean inner product between x and y .

Setup At each time $i = 1, 2, \dots, T$, one individual arrives who is characterized by a utility function $u_i : [0, 1]^k \rightarrow \mathcal{R}$. This individual is exposed to a tax vector $x_i \in [0, 1]^k$, and makes a continuous consumption decision y_i . Public revenue is given by $\langle x_i, y_i \rangle$. Private utility is given by $u_i(y_i)$ plus their consumption of a numeraire good, which has price normalized to 1 and enters utility additively. The individual consumption choice y_i costs $\langle x_i + p, y \rangle$, where p is the (exogenously given) vector of pre-tax prices. This

cost of purchasing y_i reduces the consumption of the numeraire good. The optimal individual decision is therefore given by

$$y_i = G_i(x_i) = \operatorname{argmax}_{y \in [0,1]^k} [u_i(y) - \langle x_i + p, y \rangle]. \quad (25)$$

The implied private welfare is

$$v_i(x) = v_0 + \max_{y \in [0,1]^k} [u_i(y) - \langle x + p, y \rangle],$$

where we have added a constant v_0 , chosen such that $v_i(0) = 0$; this is just a normalization to simplify notation below.

We define social welfare as a weighted sum of public revenue and private welfare, with a weight λ for the latter. Social welfare for time period i , as a function of the tax vector x , is therefore given by

$$U_i(x_i) = \underbrace{\langle x_i, y_i \rangle}_{\text{Public revenue}} + \lambda \cdot \underbrace{v_i(x_i)}_{\text{Private welfare}}. \quad (26)$$

After period i , we observe y_i and the tax vector x_i . Nothing else is observed.

Mapping demand to welfare By the envelope theorem (Milgrom and Segal, 2002),

$$\nabla_x v_i(x) = G_i(x).$$

Let $\mathcal{V}$ be the set of differentiable functions v on $[0,1]^k$ such that $\nabla_x v \in L^2$, and such that $v(0) = 0$. Consider the following operator, mapping the demand function G into the corresponding indirect utility function v .

$$\Pi(G(\cdot)) \in \operatorname{argmin}_{v(\cdot) \in \mathcal{V}} \int_{[0,1]^k} \|\nabla_x v(x) - G(x)\|^2 dx \quad (27)$$

We can think of the operator Π as combining two operators. First, the function G is projected on the subspace of functions on $[0,1]^k$ which can be written as the gradient of some function v . Second, the projected G is then integrated to get $v(x)$ for any x . Integration here is along some curve in $[0,1]^k$ from 0 to x . Given the first projection, the choice of curve does not matter for the resulting function v . A formal analysis of Tempered Exp3 for Commodity Taxation would need to prove existence of the projec-

Algorithm 4 Tempered Exp3 for Commodity Taxation

Require: Tuning parameters K , γ and η .

- 1: Calculate the set of evenly spaced grid-points $\mathcal{X} = \{0, \frac{1}{K}, \dots, 1\}^k$
and initialize $\widehat{\mathbb{G}}_1(x) = 0$ for all grid points.
- 2: **for** individual $i = 1, 2, \dots, T$ **do**
- 3: For all $x \in \mathcal{X}$, set {Estimated welfare}

$$\widehat{\mathbb{U}}_i(x) = \langle x_i, \widehat{\mathbb{G}}_i \rangle + \lambda \cdot \widehat{v}_i(x_i). \quad (28)$$

- 4: For all $x \in \mathcal{X}$, set {Assignment probabilities}

$$p_i = (1 - \gamma) \cdot \frac{\exp(\eta \cdot \widehat{\mathbb{U}}_i(x))}{\sum_{x'} \exp(\eta \cdot \widehat{\mathbb{U}}_i(x'))} + \frac{\gamma}{(K + 1)^k}. \quad (29)$$

- 5: Choose x_i at random according to the probability distribution p_i , and query y_i accordingly.
- 6: For all $x \in [0, 1]^k$, set {Estimated demand}

$$\widetilde{\mathbb{G}}_{i+1}(x) = \widehat{\mathbb{G}}_i(x) + y_i \cdot \frac{\mathbf{1}(x_i = \lfloor x \rfloor)}{p_i} \quad (30)$$

$$\widehat{v}_{i+1}(x) = \Pi(\widetilde{\mathbb{G}}_{i+1}) \quad (31)$$

$$\widehat{\mathbb{G}}_{i+1}(x) = \nabla_x \widehat{v}_{i+1}(x). \quad (32)$$

7: **end for**

tion. We leave such a formal analysis, including lower and upper regret bounds, for future research.

7 Conclusion

Possible applications The setup introduced in Section 2 was deliberately stylized, to allow for a clear exposition of the conceptual issues that arise when adaptively maximizing social welfare. The algorithm that we proposed for this setup, and the generalizations discussed later in the paper, are nonetheless directly practically relevant. They remain appropriate in economic settings that are considerably more general than the setting described by our model.

The reasons for this generality have been elucidated by the public finance literature, cf. Chetty (2009), building on the generality of the envelope theorem, cf. Milgrom and Segal (2002); Sinander (2022). By the envelope theorem, the welfare impact of

a marginal tax change on private welfare can be calculated ignoring any behavioral responses to the tax change. This holds in generalizations of our setup that allow for almost arbitrary action spaces (including discrete and continuous, multi-dimensional, and dynamic actions), and for arbitrary preference heterogeneity. The expressions for social welfare that justify our algorithms remain unchanged under such generalizations. That said, the validity of these expressions does require the absence of income effects and of externalities. If there are income effects or externalities, the algorithms need to be modified.

Our approach is motivated by applications of algorithmic decisionmaking for public policy, where a policymaker cares about welfare, but also faces a government budget constraint. Possible application domains of our algorithm include the following. In *public health and development economics*, field experiments such as [Cohen and Dupas \(2010\)](#) vary the level of a subsidy for goods such as insecticide-treated bed nets (ITNs), estimating the impact on demand. Our algorithm could be used to find the optimal subsidy level quickly and apply it to experimental participants. A term capturing positive externalities of the use of ITNs could be added to social welfare, leaving the algorithm otherwise unchanged. In *educational economics*, many studies evaluate the impact of financial aid on college enrollment ([Dynarski et al., 2023](#)). An adaptive experiment might vary the level of aid provided, where aid is conditional on college attendance and conditional on pre-determined criteria of need or merit. In such an experiment, a variant of our algorithm for optimal income taxation might be used, where the welfare weights ω are a function of need or merit, and the outcome y is college attendance. In *environmental economics*, many experiments (e.g., [Lee et al. 2020](#)) study the impact of electricity pricing on household electricity consumption. Once again, our baseline algorithm (for binary household decisions about connecting to the grid) or our algorithm for commodity taxation (for continuous household decisions about consumption levels) could be applied, to learn optimal prices, taking into account both distributional considerations and externalities.

These examples are all drawn from public policy, where there is an intrinsic concern for social welfare. This contrasts with commercial applications, where the goal is typically to maximize (directly observable) profits by monopolist pricing ([den Boer, 2015](#)), or more generally by reserve price setting in auctions ([Nedelec et al., 2022](#)). Adaptive pricing algorithms are used in applications such as online ad auctions. A concern for welfare might enter in such commercial settings if there is a participation constraint that needs to be satisfied for consumers. Suppose for example that consumers or service

providers need to first sign up for a platform, say for e-commerce or for gig work, and then repeatedly engage in transactions on this platform. To sign up in the first place, their expected welfare needs to exceed their outside option. This constraint might then enter the platform provider’s objective, in Lagrangian form, adding a term for welfare, and leading to objectives such as those maximized by our algorithms.

An alternative approach: Thompson sampling The main algorithm proposed in this paper, Tempered Exp3 for Social Welfare, is designed to perform well in the adversarial setting. In the construction of this algorithm, no probabilistic assumptions were made about the distribution of v_i . In the stochastic framework, a sampling distribution is assumed, for instance that the v_i be i.i.d. over time. The Bayesian framework completes this by assuming a prior distribution over the parameters which govern the sampling distribution.

One popular heuristic for adaptive policy choice in the Bayesian framework is Thompson sampling (Thompson, 1933; Russo et al., 2018), also known as probability matching, which assigns a policy with probability equal to the posterior probability that this policy is optimal. In our setting, Thompson sampling could be implemented as follows. First, form a posterior for the demand function $\mathbf{G}(x) = E[y|x]$, based on all the data available from previous periods $j < i$. Sample one draw $\tilde{\mathbf{G}}(\cdot)$ from this posterior. Map this draw into a draw $\tilde{\mathbf{U}}(\cdot)$ from the posterior for $\mathbf{U}(\cdot)$ via $\tilde{\mathbf{U}}(x) = x \cdot \tilde{\mathbf{G}}(x) + \lambda \cdot \int_x^1 \tilde{\mathbf{G}}(x') dx'$. Find the maximizer $x_i = \arg\max_x \tilde{\mathbf{U}}(x)$. This is the policy recommended by Thompson sampling. We conjecture that this algorithm will outperform random assignment, but will under-explore relative to the optimal algorithm. Adding further forced exploration to this algorithm might improve cumulative welfare. A formal analysis of algorithms of this type is left for future research.

A natural class of priors for $\mathbf{G}$ are Gaussian process priors (Williams and Rasmussen, 2006). If outcomes y are conditionally normal (rather than binary, as in our baseline model), then the posterior for demand is available in closed form, and the posterior mean is equal to the best linear predictor given past outcomes y_j . Furthermore, since social welfare is a linear transformation of demand, the posterior for $\mathbf{U}$ is then also linear and available in closed form. For details, see Kasy (2018).

References

- Achddou, J., Cappé, O., and Garivier, A. (2021). Fast rate learning in stochastic first price bidding. *Asian Conference on Machine Learning*.
- Agrawal, R. (1995). Sample mean based index policies by $\mathcal{O}(\log n)$ regret for the multi-armed bandit problem. *Advances in applied probability*, 27(4):1054–1078.
- Auer, P., Cesa-Bianchi, N., and Fischer, P. (2002a). Finite-time analysis of the multi-armed bandit problem. *Machine learning*, 47:235–256.
- Auer, P., Cesa-Bianchi, N., Freund, Y., and Schapire, R. E. (2002b). The nonstochastic multiarmed bandit problem. *SIAM journal on computing*, 32(1):48–77.
- Bachoc, F., Cesari, T., Colomboni, R., and Paudice, A. (2022a). A near-optimal algorithm for univariate zeroth-order budget convex optimization.
- Bachoc, F., Cesari, T., Colomboni, R., and Paudice, A. (2022b). Regret analysis of dyadic search. *arXiv preprint arXiv:2209.00885*.
- Baily, M. (1978). Some aspects of optimal unemployment insurance. *Journal of Public Economics*, 10(3):379–402.
- Bubeck, S. and Cesa-Bianchi, N. (2012). Regret Analysis of Stochastic and Nonstochastic Multi-armed Bandit Problems. *Foundations and Trends® in Machine Learning*, 5(1):1–122.
- Cesa-Bianchi, N., Cesari, T., Colomboni, R., Fusco, F., and Leonardi, S. (2024a). Bilateral trade: A regret minimization perspective. *Mathematics of Operations Research*, 49(1):171–203.
- Cesa-Bianchi, N., Cesari, T., Colomboni, R., Fusco, F., and Leonardi, S. (2024b). The role of transparency in repeated first-price auctions with unknown valuations. In *Proceedings of the 56th Annual ACM Symposium on Theory of Computing*, pages 225–236.
- Cesa-Bianchi, N., Cesari, T. R., Colomboni, R., Fusco, F., and Leonardi, S. (2021). A regret analysis of bilateral trade. In *Proceedings of the 22nd ACM Conference on Economics and Computation*, pages 289–309.

- Cesa-Bianchi, N., Cesari, T. R., Colomboni, R., Fusco, F., and Leonardi, S. (2023). Repeated bilateral trade against a smoothed adversary. *The Thirty Sixth Annual Conference on Learning Theory*, pages 1095–1130.
- Cesa-Bianchi, N., Freund, Y., Haussler, D., Helmbold, D. P., Schapire, R. E., and Warmuth, M. K. (1997). How to use expert advice. *Journal of the ACM (JACM)*, 44(3):427–485.
- Cesa-Bianchi, N., Gaillard, P., Gentile, C., and Gerchinovitz, S. (2017). Algorithmic chaining and the role of partial feedback in online nonparametric learning. In *Conference on Learning Theory*, pages 465–481. PMLR.
- Cesa-Bianchi, N., Gentile, C., and Mansour, Y. (2015). Regret minimization for reserve prices in second-price auctions. *IEEE Transactions on Information Theory*, 1(61):549–564.
- Cesa-Bianchi, N. and Lugosi, G. (2006). *Prediction, learning, and games*. Cambridge University Press.
- Chetty, R. (2009). Sufficient statistics for welfare analysis: A bridge between structural and reduced-form methods. *Annual Review of Economics*, 1(1):451–488.
- Cohen, J. and Dupas, P. (2010). Free distribution or cost-sharing? evidence from a randomized malaria prevention experiment. *The Quarterly Journal of Economics*, 125(1):1–45.
- Cover, T. (1965). Behavior of sequential predictors of binary sequences. In *Proceedings of the 4th Prague Conference on Information Theory, Statistical Decision Functions, Random Processes*, pages 263–272. Publishing House of the Czechoslovak Academy of Sciences, Prague.
- Cover, T. M. and Thomas, J. A. (2006). *Elements of Information Theory*. Wiley-Interscience.
- Daskalakis, C. and Syrgkanis, V. (2022). Learning in auctions: Regret is hard, envy is easy. *Games and Economic Behavior*.
- den Boer, A. V. (2015). Dynamic pricing and learning: historical origins, current research, and new directions. *Surveys in Operations Research and Management Science*.

- Dynarski, S., Page, L., and Scott-Clayton, J. (2023). College costs, financial aid, and student decisions. In *Handbook of the Economics of Education*, volume 7, pages 227–285. Elsevier.
- Feng, Z., Guruganesh, G., Liaw, C., Mehta, A., and Sethi, A. (2021). Convergence analysis of no-regret bidding algorithms in repeated auctions. In *Proceedings of the AAAI Conference on Artificial Intelligence*, pages 5399–5406.
- Feng, Z., Podimata, C., and Syrgkanis, V. (2018). Learning to bid without knowing your value. In *Proceedings of the 2018 ACM Conference on Economics and Computation*, pages 505–522.
- Fudenberg, D. and Levine, D. K. (1995). Consistency and cautious fictitious play. *Journal of Economic Dynamics and Control*, 19(5-7):1065–1089.
- Han, Y., Zhou, Z., Flores, A., Ordentlich, E., and Weissman, T. (2020a). Learning to bid optimally and efficiently in adversarial first-price auctions. *arXiv preprint arXiv:2007.04568*.
- Han, Y., Zhou, Z., and Weissman, T. (2020b). Optimal no-regret learning in repeated first-price auctions. *arXiv preprint arXiv:2003.09795*.
- Hannan, J. (1957). Approximation to bayes risk in repeated play. *Contributions to the Theory of Games*, 3(2):97–139.
- Hart, S. and Mas-Colell, A. (2000). A simple adaptive procedure leading to correlated equilibrium. *Econometrica*, 68(5):1127–1150.
- Hart, S. and Mas-Colell, A. (2001). A general class of adaptive strategies. *Journal of Economic Theory*, 98(1):26–54.
- Kasy, M. (2018). Optimal taxation and insurance using machine learning – sufficient statistics and beyond. *Journal of Public Economics*, 167.
- Kasy, M. and Sautmann, A. (2021). Adaptive treatment assignment in experiments for policy choice. *Econometrica*, 89(1):113–132.
- Kleinberg, R. D. and Leighton, F. T. (2003). The value of knowing a demand curve: Bounds on regret for online posted-price auctions. In *IEEE Symposium on Foundations of Computer Science*, pages 594–605.

- Kolumbus, Y. and Nisan, N. (2022). Auctions between regret-minimizing agents. In *Proceedings of the ACM Web Conference 2022*, pages 100–111.
- Lai, T. L. (1987). Adaptive treatment allocation and the multi-armed bandit problem. *The annals of statistics*, pages 1091–1114.
- Lai, T. L. and Robbins, H. (1985). Asymptotically efficient adaptive allocation rules. *Advances in applied mathematics*, 6(1):4–22.
- Lattimore, T. and Szepesvári, C. (2020). *Bandit algorithms*. Cambridge University Press.
- Lee, K., Miguel, E., and Wolfram, C. (2020). Experimental evidence on the economics of rural electrification. *Journal of Political Economy*, 128(4):1523–1565.
- Littlestone, N. and Warmuth, M. (1994). The weighted majority algorithm. *Information and computation*, 108(2):212–261.
- Lucier, B., Pattathil, S., Slivkins, A., and Zhang, M. (2024). Autobidders with budget and roi constraints: Efficiency, regret, and pacing dynamics. In *The Thirty Seventh Annual Conference on Learning Theory*, pages 3642–3643. PMLR.
- Lugosi, G., Markakis, M. G., and Neu, G. (2023). On the hardness of learning from censored and nonstationary demand. *INFORMS Journal on Optimization*.
- Milgrom, P. and Segal, I. (2002). Envelope theorems for arbitrary choice sets. *Econometrica*, 70(2):583–601.
- Mirrlees, J. (1971). An exploration in the theory of optimum income taxation. *The Review of Economic Studies*, pages 175–208.
- Nedelec, T., Calauzènes, C., El Karoui, N., Perchet, V., et al. (2022). Learning in repeated auctions. *Foundations and Trends® in Machine Learning*, 15(3):176–334.
- Ramsey, F. P. (1927). A contribution to the theory of taxation. *The economic journal*, 37(145):47–61.
- Russo, D. (2020). Simple bayesian algorithms for best-arm identification. *Operations Research*, 68(6):1625–1647.

- Russo, D. J., Roy, B. V., Kazerouni, A., Osband, I., and Wen, Z. (2018). A Tutorial on Thompson Sampling. *Foundations and Trends® in Machine Learning*, 11(1):1–96.
- Saez, E. (2001). Using elasticities to derive optimal income tax rates. *The Review of Economic Studies*, 68(1):205–229.
- Saez, E. (2002). Optimal income transfer programs: intensive versus extensive labor supply responses. *The Quarterly Journal of Economics*, 117(3):1039–1073.
- Saez, E. and Stantcheva, S. (2016). Generalized social welfare weights for optimal tax theory. *American Economic Review*, 106(1):24–45.
- Seldin, Y. and Slivkins, A. (2014). One practical algorithm for both stochastic and adversarial bandits. In *International Conference on Machine Learning*, pages 1287–1295. PMLR.
- Sinander, L. (2022). The converse envelope theorem. *Econometrica*, 90(6):2795–2819.
- Slivkins, A. (2019). Introduction to multi-armed bandits. *arXiv preprint arXiv:1904.07272*.
- Thompson, W. R. (1933). On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 25(3/4):285–294.
- Tsybakov, A. B. (2008). *Introduction to Nonparametric Estimation*. Springer Publishing Company, Incorporated.
- Vovk, V. G. (1990). Aggregating strategies. In *Proceedings of the 3rd Annual Workshop on Computational Learning Theory*, pages 371–386.
- Weed, J., Perchet, V., and Rigollet, P. (2016). Online learning in repeated auctions. In *Conference on Learning Theory*, pages 1562–1583. PMLR.
- Williams, C. and Rasmussen, C. (2006). *Gaussian processes for machine learning*. MIT Press.
- Williams, D. (1991). *Probability with martingales*. Cambridge University Press.
- Zimmert, J. and Seldin, Y. (2021). Tsallis-inf: An optimal algorithm for stochastic and adversarial bandits. *Journal of Machine Learning Research*, 22(28):1–49.

A Proofs

A.1 Theorem 1 (Lower bound on regret)

Defining a family of distributions for v Recall that, for each $\epsilon \in [-1, 1]$, the probability distribution μ^ϵ is defined as the probability measure supported on $(1/4, 1/2, 3/4, 1)$ with masses $(a, (1 + \epsilon) \cdot b, (1 - \epsilon) \cdot b, 1 - a - 2 \cdot b)$, where

$$a := \frac{(1 - \lambda) \cdot (136 - 99 \cdot \lambda)}{2 \cdot (4 - 3 \cdot \lambda) \cdot (24 - 17 \cdot \lambda)}, \quad b := \frac{1 - \lambda}{2 \cdot (24 - 17 \cdot \lambda)}.$$

Furthermore, for each $\epsilon \in [-1, 1]$, recall that $\mathbf{G}^\epsilon$ and $\mathbf{U}^\epsilon$ are respectively the demand function and the expected social welfare associated to μ^ϵ (see Figure 1 for an illustration). Let $v_1, v_2, \dots \in [0, 1]$ be the sequence of individual valuations. For each $\epsilon \in [-1, 1]$, consider a distribution P^ϵ such that the individual valuations $v_1, v_2, \dots$ form a P^ϵ -i.i.d. sequence (independent of the randomization used by the algorithm) with common distribution μ^ϵ .

Explicit lower bound on regret that will be proven Define

$$c_1 := \frac{\lambda}{4} \cdot b, \quad c_2 := \frac{1}{8} \cdot \frac{1 - \lambda}{4 - 3 \cdot \lambda}, \quad c_3 := b \cdot \sqrt{\frac{2}{a \cdot (1 - a - 2 \cdot b)}}.$$

We will prove that, for any randomized algorithm and any time horizon $T \in \mathbb{N}$, there exists $\epsilon \in [-1, 1]$ such that

$$\mathcal{R}_T(\mathbf{G}^\epsilon) \geq C \cdot T^{2/3},$$

where

$$\begin{aligned} C &:= \min \left(\frac{c_1^2 \cdot c_3^2}{c_2}, \frac{c_2}{2}, \frac{1}{16} \cdot \sqrt[3]{\frac{c_1^2 \cdot c_2}{c_3^2}} \right) \\ &= \min \left(\frac{\lambda^2 \cdot (4 - 3 \cdot \lambda)^3}{8 \cdot (136 - 99 \cdot \lambda) \cdot (26 - 19 \cdot \lambda)}, \frac{\lambda^{2/3} \cdot (1 - \lambda)^{4/3} \cdot (136 - 99 \cdot \lambda)^{1/3} \cdot (26 - 19 \cdot \lambda)^{1/3}}{128 \cdot (4 - 3 \cdot \lambda) \cdot (24 - 17 \cdot \lambda)^{4/3}} \right) > 0 \end{aligned} \tag{33}$$

Fix a randomized algorithm to choose the policies $x_1, x_2, \dots$, and fix a time horizon $T \in \mathbb{N}$.

Number of mistakes and lower bound on regret We need to count the random number of times the algorithm has played in the regions $(1/2, 3/4]$, $[0, 1/2]$ and $(3/4, 1]$ up to time T . This can be done relying on the following random variables:

$$n_1 := \sum_{i=1}^T \mathbf{1}_{(1/2, 3/4]}(x_i) , \quad n_2 := \sum_{i=1}^T \mathbf{1}_{[0, 1/2]}(x_i) , \quad n_3 := \sum_{i=1}^T \mathbf{1}_{(3/4, 1]}(x_i) .$$

Notice that since the intervals $(1/2, 3/4]$, $[0, 1/2]$ and $(3/4, 1]$ form a partition of $[0, 1]$, we have that

$$n_1 + n_2 + n_3 = T \tag{34}$$

For each $\epsilon \in [-1, 1]$, denote by E^ϵ the expectation taken with respect to the distribution P^ϵ . Notice that, for each $\epsilon \in [-1, 1]$, the expected regret when the underlying distribution is P^ϵ equals

$$\mathcal{R}_T(\mathbf{G}^\epsilon) = T \cdot \sup_{x \in [0, 1]} \mathbf{U}^\epsilon(x) - \sum_{i=1}^T E^\epsilon(\mathbf{U}^\epsilon(x_i)) . \tag{35}$$

Algebraic calculations show that, for each $\epsilon \in [-1, 1]$

$$\max_{x \in (1/2, 3/4]} \mathbf{U}^\epsilon(x) = \mathbf{U}^\epsilon(3/4) , \quad \max_{x \in [0, 1/2]} \mathbf{U}^\epsilon(x) = \mathbf{U}^\epsilon(1/4) , \quad \max_{x \in (3/4, 1]} \mathbf{U}^\epsilon(x) = \mathbf{U}^\epsilon(1) , \tag{36}$$

$$\text{and} \quad \mathbf{U}^\epsilon(1) - \mathbf{U}^\epsilon(1/4) = c_1 \cdot \epsilon . \tag{37}$$

Further calculations show also that

$$\min_{\epsilon \in [-1, 1]} \min(\mathbf{U}^\epsilon(1/4), \mathbf{U}^\epsilon(1)) = \mathbf{U}^1(1/4) , \quad \max_{\epsilon \in [-1, 1]} \max_{x \in (1/2, 3/4]} \mathbf{U}^\epsilon(x) = \mathbf{U}^{-1}(3/4) , \tag{38}$$

$$\text{and} \quad \mathbf{U}^1(1/4) - \mathbf{U}^{-1}(3/4) = c_2 . \tag{39}$$

Equations (36), (37), (38), and (39) imply that

$$\sup_{x \in [0, 1]} \mathbf{U}^\epsilon(x) = \mathbf{U}^\epsilon(1) , \quad \text{if } \epsilon \in [0, 1] . \tag{40}$$

It follows that, if $\epsilon \in [0, 1]$,

$$\begin{aligned}
\mathcal{R}_T(\mathbf{G}^\epsilon) &\stackrel{(35)}{=} T \cdot \sup_{x \in [0,1]} U^\epsilon(x) - \sum_{i=1}^T E^\epsilon(U^\epsilon(x_i)) \\
&\stackrel{(40)}{=} T \cdot U^\epsilon(1) - \sum_{i=1}^T E^\epsilon\left(U^\epsilon(x_i) \cdot (\mathbf{1}_{(1/2, 3/4]}(x_i) + \mathbf{1}_{[0, 1/2]}(x_i) + \mathbf{1}_{(3/4, 1]}(x_i))\right) \\
&\stackrel{(36)}{\geq} T \cdot U^\epsilon(1) - \sum_{i=1}^T E^\epsilon\left(U^\epsilon(3/4) \cdot \mathbf{1}_{(1/2, 3/4]}(x_i) + U^\epsilon(1/2) \cdot \mathbf{1}_{[0, 1/2]}(x_i) + U^\epsilon(1) \cdot \mathbf{1}_{(3/4, 1]}(x_i)\right) \\
&\stackrel{(34)}{=} (U^\epsilon(1) - U^\epsilon(3/4)) \cdot E^\epsilon(n_1) + (U^\epsilon(1) - U^\epsilon(1/4)) \cdot E^\epsilon(n_2) \\
&\stackrel{(38)}{\geq} (U^1(1/4) - U^{-1}(3/4)) \cdot E^\epsilon(n_1) + (U^\epsilon(1) - U^\epsilon(1/4)) \cdot E^\epsilon(n_2) \\
&\stackrel{(39)}{=} c_2 \cdot E^\epsilon(n_1) + (U^\epsilon(1) - U^\epsilon(1/4)) \cdot E^\epsilon(n_2) \\
&\stackrel{(37)}{=} c_2 \cdot E^\epsilon(n_1) + c_1 \cdot \epsilon \cdot E^\epsilon(n_2) \tag{41}
\end{aligned}$$

Notice that inequality (41) quantifies how much regret the algorithm is going to suffer in terms of the expected number of times it plays in the wrong regions, when the demand function is $\mathbf{G}^\epsilon$ and $\epsilon > 0$.

In the same way inequality (41) was proven, we can prove that, if $\epsilon \in [0, 1]$,

$$\mathcal{R}_T(\mathbf{G}^{-\epsilon}) \geq c_2 \cdot E^{-\epsilon}(n_1) + c_1 \cdot \epsilon \cdot E^{-\epsilon}(n_3) \geq c_1 \cdot \epsilon \cdot E^{-\epsilon}(n_3) , \tag{42}$$

which again quantifies how much regret the algorithm is going to suffer in terms of the expected number of times it plays in the wrong regions, when the demand function is $\mathbf{G}^{-\epsilon}$ and $\epsilon > 0$.

Intuition for the remainder of the proof At high level, inequalities (41) and (42) tell us that, if $|\epsilon|$ is not negligible, the algorithm has to play a substantially different number of times in the region $(3/4, 1]$, depending on the sign of ϵ , not to suffer significant regret when the demand function is $\mathbf{G}^\epsilon$. The crucial idea is that the only way for the algorithm to present this different behavior is by playing in the only informative region about the sign of ϵ , i.e., the region $(1/2, 3/4]$. However, as shown in (41), selecting policies in this region comes at a cost in terms of regret. To relate quantitatively the number of times the algorithm has to play in this costly region with the difference in the expected number of times the algorithm selects policies in the region $(3/4, 1]$ is the last missing

ingredient that we can obtain relying on information theoretic techniques: It can be proved (and a formal proof is provided in the online appendix, in Section [B.1](#)) that, for each $\epsilon \in [0, 1]$,

$$E^{-\epsilon}(n_3) \geq E^\epsilon(n_3) - c_3 \cdot \epsilon \cdot T \cdot \sqrt{E^\epsilon(n_1)}. \quad (43)$$

Now, if the algorithm is going to suffer low regret when $\epsilon > 0$, then by [\(41\)](#) we have an upper bound on the number of times the algorithm plays in the region $(1/2, 3/4]$ and a lower bound on the number of times it plays in the region $(3/4, 1]$, whenever $\epsilon > 0$. In turn, by [\(43\)](#), this gives a lower bound on the number of times the algorithm plays in the sub-optimal region $(3/4, 1]$ when $\epsilon < 0$. Then, relying on [\(42\)](#), we have an explicit lower bound on how much regret the algorithm is going to suffer when $\epsilon < 0$. We will now carry out this plan —and prove the theorem— as follows.

Low regret cannot be achieved for both positive and negative ϵ To get a contradiction, suppose that

$$\forall \epsilon \in [-1, 1] \quad \mathcal{R}_T(\mathbf{G}^\epsilon) < C \cdot T^{2/3}. \quad (44)$$

It follows from [\(41\)](#) that, for each $\epsilon \in [0, 1]$,

$$E^\epsilon(n_1) \stackrel{(41)}{\leq} \frac{\mathcal{R}_T(\mathbf{G}^\epsilon)}{c_2} \stackrel{(44)}{\leq} \frac{C}{c_2} \cdot T^{2/3}, \quad E^\epsilon(n_2) \stackrel{(41)}{\leq} \frac{\mathcal{R}_T(\mathbf{G}^\epsilon)}{c_1 \cdot \epsilon} \stackrel{(44)}{\leq} \frac{C}{c_1 \cdot \epsilon} \cdot T^{2/3}. \quad (45)$$

This implies, relying also on [\(42\)](#) and [\(43\)](#), that for each $\epsilon \in [0, 1]$ we have

$$\begin{aligned} \mathcal{R}_T(\mathbf{G}^{-\epsilon}) &\stackrel{(42)}{\geq} c_1 \cdot \epsilon \cdot E^{-\epsilon}(n_3) \stackrel{(43)}{\geq} c_1 \cdot \epsilon \cdot (E^\epsilon(n_3) - c_3 \cdot \epsilon \cdot T \cdot \sqrt{E^\epsilon(n_1)}) \\ &\stackrel{(34)}{=} c_1 \cdot \epsilon \cdot (T - E^\epsilon(n_1) - E^\epsilon(n_2) - c_3 \cdot \epsilon \cdot T \cdot \sqrt{E^\epsilon(n_1)}) \\ &\stackrel{(45)}{\geq} c_1 \cdot \epsilon \cdot \left(T - \frac{C}{c_2} \cdot T^{2/3} - \frac{C}{c_1 \cdot \epsilon} \cdot T^{2/3} - c_3 \cdot \epsilon \cdot T \cdot \sqrt{\frac{C}{c_2} \cdot T^{2/3}} \right) \\ &= c_1 \cdot \epsilon \cdot \left(1 - \frac{C}{c_2} \cdot T^{-1/3} - \frac{C}{c_1 \cdot \epsilon} \cdot T^{-1/3} - c_3 \cdot \epsilon \cdot T^{1/3} \cdot \sqrt{\frac{C}{c_2}} \right) \cdot T. \end{aligned} \quad (46)$$

Pick $\epsilon := T^{-1/3} \cdot \sqrt{\frac{\sqrt{C \cdot c_2}}{c_1 \cdot c_3}}$. First, note that since $0 < C \stackrel{(33)}{\leq} \frac{c_1^2 \cdot c_3^2}{c_2}$ we have that $\epsilon \in (0, 1]$. Plugging this value of ϵ in (46) leads to

$$\begin{aligned}
C \cdot T^{2/3} &\stackrel{(44)}{>} \mathcal{R}_T(\mathbf{G}^{-\epsilon}) \\
&\stackrel{(46)}{\geq} \sqrt{\frac{\sqrt{C \cdot c_2} \cdot c_1}{c_3}} \cdot \left(1 - \frac{C}{c_2} \cdot T^{-1/3} - 2 \cdot \sqrt{\frac{c_3}{c_1 \cdot \sqrt{c_2}}} \cdot C^{3/4}\right) \cdot T^{2/3} \\
&\stackrel{(33)}{\geq} \frac{1}{2} \cdot \sqrt{\frac{\sqrt{C \cdot c_2} \cdot c_1}{c_3}} \cdot \left(1 - 4 \cdot \sqrt{\frac{c_3}{c_1 \cdot \sqrt{c_2}}} \cdot C^{3/4}\right) \cdot T^{2/3} \\
&\stackrel{(33)}{\geq} \frac{1}{4} \cdot \sqrt{\frac{\sqrt{C \cdot c_2} \cdot c_1}{c_3}} \cdot T^{2/3}, \tag{47}
\end{aligned}$$

where the second to last inequality follows from $C \leq \frac{c_2}{2}$, while the last inequality follows from $C \leq \frac{1}{16} \sqrt[3]{\frac{c_1^2 \cdot c_2}{c_3^2}}$. Rearranging inequality (47) leads to the contradiction

$$C \stackrel{(47)}{>} \left(\frac{1}{4} \cdot \sqrt{\frac{c_1 \cdot \sqrt{c_2}}{c_3}}\right)^{4/3} = \frac{1}{8} \cdot \sqrt[3]{\frac{2 \cdot c_1^2 \cdot c_2}{c_3^2}} > \frac{1}{16} \cdot \sqrt[3]{\frac{c_1^2 \cdot c_2}{c_3^2}} \stackrel{(33)}{\geq} C.$$

Since (44) leads to a contradiction, it follows that there exists $\epsilon \in [-1, 1]$ such that $\mathcal{R}_T(\mathbf{G}^\epsilon) \geq C \cdot T^{2/3}$. Given that the time horizon T and the randomized algorithm were arbitrarily fixed, the theorem is proved.

A.2 Theorem 2 (Adversarial upper bound on regret)

The proof of this theorem builds upon the proof of Theorem 6.5 in [Cesa-Bianchi and Lugosi \(2006\)](#). Relative to this theorem, we need to additionally consider the discretization error introduced by Algorithm 1, and explicitly control the variance of estimated welfare.

Recall our notation $\mathbb{U}$ and $\mathbb{U}(x)$ for realized cumulative welfare, and for cumulative welfare for the counterfactual, fixed policy x . We further abbreviate $\mathbb{U}_{T_k} = \mathbb{U}(\tilde{x}_k)$. Throughout this proof, the sequence $\{v_i\}_{i=1}^T$ is given and conditioned on in any expectations.

1. Discretization

Recall that $U_i(x) = x \cdot \mathbf{1}(x \leq v_i) + \lambda \cdot \max(v_i - x, 0)$. Let

$$\tilde{v}_i = \max_k \{\tilde{x}_k : \tilde{x}_k \leq v_i\}$$

(this is v_i rounded down to the next gridpoint $\tilde{x}_k$), and denote

$$\begin{aligned}\tilde{U}_i(x) &= x \cdot \mathbf{1}(x \leq v_i) + \lambda \cdot \max(\tilde{v}_i - x, 0), \\ \tilde{\mathbb{U}}_i(x) &= \sum_{j \leq i} \tilde{U}_j(x),\end{aligned}$$

as well as $\tilde{\mathbb{U}}_{ik} = \tilde{\mathbb{U}}_i(\tilde{x}_k)$. Then it is immediate that $\tilde{U}_i(x) \leq U_i(x)$,

$$\sup_x |\tilde{U}_i(x) - U_i(x)| \leq \frac{\lambda}{K},$$

and $\operatorname{argmax}_x \tilde{\mathbb{U}}_i(x) \in \{\tilde{x}_1, \dots, x_{K+1}\}$, and therefore

$$\max_k \tilde{\mathbb{U}}_{ik} \geq \sup_x \mathbb{U}_i(x) - i \cdot \frac{\lambda}{K}$$

2. Unbiasedness

At the end of period i , $\hat{G}_k$ is an unbiased estimator of $\sum_{j \leq i} \mathbf{1}(\tilde{x}_k \leq v_j)$ for all k . Therefore, $E[\hat{\mathbb{U}}_{ik}] = \tilde{\mathbb{U}}_{ik}$ for all i and k .

3. Upper bound on optimal welfare

Define $W_i = \sum_k \exp(\eta \cdot \hat{\mathbb{U}}_{ik})$, and $q_{ik} = \exp(\eta \cdot \hat{\mathbb{U}}_{ik})/W_i$.

It is immediate that,

$$E[\log W_T] \geq \eta \cdot E[\max_k \hat{\mathbb{U}}_{Tk}] \geq \eta \cdot \max_k E[\hat{\mathbb{U}}_{Tk}] = \eta \cdot \max_k \tilde{\mathbb{U}}_{Tk}.$$

Furthermore

$$E[\log W_T] = \sum_{0 \leq i < T} E \left[\log \left(\frac{W_{i+1}}{W_i} \right) \right] + \log(W_0).$$

Given our initialization of the algorithm, $\log(W_0) = \log(K + 1)$.

4. Lower bound on estimated welfare

Denote $\hat{U}_{ik} = \tilde{x}_k \cdot \hat{H}_k + \frac{\lambda}{K} \cdot \sum_{k' > k} \hat{H}_{k'}$, where $\hat{H}_k = \frac{y_i}{p_{ik}} \cdot \mathbf{1}(k_i = k)$,

so that $\widehat{U}_{ik} = \sum_{j < i} \widehat{U}_{jk}$, and $E[\widehat{U}_{jk}] = U_i(\tilde{x}_k)$.

By definition of W_i ,

$$\log \left(\frac{W_{i+1}}{W_i} \right) = \log \left(\sum_k q_{ik} \cdot \exp(\eta \cdot \widehat{U}_{ik}) \right).$$

Since $p_k \geq \gamma/(K+1)$ for all k , $\widehat{U}_{ik} \in [0, (K+1)/\gamma]$ for all i and k , and therefore $\eta \cdot \widehat{U}_{ik} \leq (K+1) \cdot \eta/\gamma \leq 1$ (where the last inequality holds by assumption). Using $\exp(a) \leq 1 + a + (e-2)a^2$ for any $a \leq 1$ yields

$$\exp(\eta \widehat{U}_{ik}) \leq 1 + \eta \cdot \widehat{U}_{ik} + (e-2) \cdot (\eta \cdot \widehat{U}_{ik})^2.$$

Therefore,

$$\begin{aligned} \log \left(\frac{W_{i+1}}{W_i} \right) &\leq \log \left(\sum_k q_{ik} \cdot \left(1 + \eta \cdot \widehat{U}_{ik} + (e-2) \cdot (\eta \cdot \widehat{U}_{ik})^2 \right) \right) \\ &\leq \eta \cdot \sum_k q_{ik} \cdot \widehat{U}_{ik} + (e-2) \cdot \eta^2 \cdot \sum_k q_{ik} \cdot \widehat{U}_{ik}^2 \end{aligned}$$

The second inequality follows from $\log(1+x) \leq x$.

5. Connecting the first order term to welfare

Note that, by definition, $q_{ik} = (p_{ik} - \frac{\gamma}{K+1}) / (1-\gamma)$. Therefore

$$\sum_k q_{ik} \cdot \widehat{U}_{ik} = \frac{1}{1-\gamma} \sum_k p_{ik} \cdot \widehat{U}_{ik} - \frac{\gamma}{(1-\gamma)(K+1)} \cdot \sum_k \widehat{U}_{ik},$$

and thus

$$E \left[\sum_k q_{ik} \cdot \widehat{U}_{ik} \right] \leq \frac{1}{1-\gamma} E \left[\tilde{U}_i(x_i) \right],$$

where we have used the fact that $0 \leq \tilde{U}_k \leq 1$ for all k , given our definition of $\tilde{U}$, and the fact that k_i is distributed according to p_{ik} , by construction.

6. Bounding the second moment of estimated welfare

It remains to bound the term $E \left[\sum_k q_{ik} \cdot \widehat{U}_{ik}^2 \right]$. As in the preceding item, we have

$$\sum_k q_{ik} \cdot \widehat{U}_{ik}^2 \leq \frac{1}{1-\gamma} \sum_k p_{ik} \cdot \widehat{U}_{ik}^2.$$

We can rewrite

$$\widehat{U}_{ik} = (\tilde{x}_k \cdot \mathbf{1}(k_i = k) + \frac{\lambda}{K} \cdot \mathbf{1}(k_i > k)) \cdot \frac{y_i}{p_{ik_i}}.$$

Bounding $y_i \leq 1$ immediately gives

$$E_i \left[\widehat{U}_{ik}^2 \right] \leq \frac{\tilde{x}_k^2}{p_{ik}} + \left(\frac{\lambda}{K} \right)^2 \cdot \sum_{k' > k} \frac{1}{p_{ik'}},$$

and therefore

$$\begin{aligned} E_i \left[\sum_k p_{ik} \cdot \widehat{U}_{ik}^2 \right] &\leq \sum_k \tilde{x}_k^2 + \left(\frac{\lambda}{K} \right)^2 \cdot \sum_k \sum_{k' > k} \frac{p_{ik}}{p_{ik'}} \\ &\leq \sum_k \left(\frac{k}{K} \right)^2 + \left(\frac{\lambda}{K} \right)^2 \cdot \sum_k p_{ik} \sum_{k' \neq k} \frac{K+1}{\gamma} \\ &= \frac{K(K+1)(2K+1)}{6K^2} + \frac{\lambda^2}{\gamma} \frac{K+1}{K} \\ &= \frac{K+1}{K} \cdot \left(\frac{2K+1}{6} + \frac{\lambda^2}{\gamma} \right). \end{aligned}$$

7. Collecting inequalities

Combining the preceding items, we get

$$\begin{aligned} &\eta \cdot \left(\sup_x \mathbb{U}(x) - T \cdot \frac{\lambda}{K} \right) \\ &\leq \eta \cdot \max_k \tilde{\mathbb{U}}_{Tk} \leq E[\log W_T] \quad (\text{Item 1}) \\ &= \sum_{0 \leq i < T} E \left[\log \left(\frac{W_{i+1}}{W_i} \right) \right] + \log(K+1) \quad (\text{Item 3}) \\ &\leq \frac{\eta}{1-\gamma} \cdot E \left[\tilde{\mathbb{U}} \right] + (e-2) \cdot \frac{\eta^2}{1-\gamma} \sum_{1 \leq i \leq T} \sum_k E \left[p_{ik} \cdot \widehat{U}_{ik}^2 \right] + \log(K+1) \quad (\text{Item 4 and 5}) \\ &\leq \frac{\eta}{1-\gamma} \cdot E \left[\tilde{\mathbb{U}} \right] + (e-2) \cdot \frac{\eta^2}{1-\gamma} T \cdot \frac{K+1}{K} \cdot \left(\frac{2K+1}{6} + \frac{\lambda^2}{\gamma} \right) + \log(K+1). \quad (\text{Item 6}) \end{aligned}$$

Multiplying by $(1 - \gamma)$ and dividing by η , adding $\gamma \sup_x \mathbb{U}(x) + T \frac{\lambda}{K}$ to both sides and subtracting $E [\tilde{\mathbb{U}}]$, bounding $\sup_x \mathbb{U}(x) \leq T$, and $E [\tilde{\mathbb{U}}] \leq E [\mathbb{U}]$ (from Item 1), yields

$$\begin{aligned} & \sup_x \mathbb{U}(x) - E [\mathbb{U}] \\ & \leq \left(\gamma + \eta \cdot (e - 2) \frac{K+1}{K} \cdot \left(\frac{2K+1}{6} + \frac{\lambda^2}{\gamma} \right) + \frac{\lambda}{K} \right) \cdot T + \frac{\log(K+1)}{\eta}. \end{aligned} \quad (48)$$

This proves the first claim of the theorem.

8. Optimizing tuning parameters

Suppose now that we choose the tuning parameters as follows:

$$\gamma = c_1 \cdot \left(\frac{\log(T)}{T} \right)^{1/3}, \quad \eta = c_2 \cdot \gamma^2, \quad K = c_3 / \gamma.$$

Plugging in we get

$$\begin{aligned} & \sup_x \mathbb{U}(x) - E [\mathbb{U}] \\ & \leq \left(\gamma + c_2 \cdot \gamma^2 \cdot (e - 2) \frac{K+1}{K} \cdot \left(\frac{2c_3/\gamma+1}{6} + \frac{\lambda^2}{\gamma} \right) + \lambda \cdot \gamma / c_3 \right) \cdot T + \frac{\log(K+1)}{c_2 \cdot \gamma^2} \\ & = \log(T)^{1/3} T^{2/3} \cdot \left(c_1 + (e - 2) \frac{K+1}{K} \cdot c_1 c_2 \left(\frac{c_3}{3} + \lambda^2 + \frac{\gamma}{6} \right) + \lambda \frac{c_1}{c_3} + \frac{\log(T^{1/3} \log(T)^{-1/3} c_3 / c_1 + 1)}{c_1^2 \log(T)} \right) \\ & = \log(T)^{1/3} T^{2/3} \cdot \left(c_1 + (e - 2) \cdot c_1 c_2 \left(\frac{c_3}{3} + \lambda^2 \right) + \lambda \frac{c_1}{c_3} + \frac{1}{3c_1^2} + o(1) \right). \end{aligned}$$

The second claim of the theorem follows.

Online Appendix

B Additional proofs

Contents

B.1 Claim (43) (Relating choice probabilities for positive and negative ϵ)	45
B.2 Theorem 3 (Lower bound on regret for the concave case)	50
B.3 Theorem 4 (Stochastic upper bound on regret of Dyadic Search for Social Welfare)	54
B.3.1 Lemma 1 (Confidence intervals contain true welfare differences with high probability)	56
B.3.2 Lemma 2 (Confidence intervals shrink with epoch length)	57
B.3.3 Theorem 4 (Completing the proof)	60
B.4 Theorem 5 (Upper bound on regret of Tempered Exp3 for Optimal Income Taxation)	65

B.1 Claim (43) (Relating choice probabilities for positive and negative ϵ)

Proof of Claim (43).

Let $w_1, w_2, \dots \in [0, 1]$ be the randomization seeds to be used by the algorithm. In the light of the Skorokhod representation theorem (Williams, 1991, Section 17.3), we may assume without (much) loss of generality that, for each $\epsilon \in [-1, 1]$, these seeds form a sequence of P^ϵ -i.i.d. $[0, 1]$ -valued uniform random variables. In particular, this implies,

$$P_{(w_i)_{i \in \mathbb{N}}}^\epsilon = P_{(w_i)_{i \in \mathbb{N}}}^{-\epsilon}, \quad \forall \epsilon \in [0, 1]. \quad (49)$$

Recall that a sequence of functions $\alpha := (\alpha_i)_{i \in \mathbb{N}}$ is called a randomized algorithm if

$$\alpha_1: [0, 1] \rightarrow [0, 1], \quad \forall i \in \mathbb{N}, \quad \alpha_{i+1}: [0, 1]^{i+1} \times \{0, 1\}^i \rightarrow [0, 1].$$

The feedback function associated to our problem is

$$\varphi: [0, 1] \times \{1/4, 1/2, 3/4, 1\} \rightarrow \{0, 1\}, \quad (x, v) \mapsto \mathbf{1}(x \leq v).$$

Now, a randomized algorithm α generates a sequence of choices $x_1, x_2, \dots$ using the randomization seeds $w_1, w_2, \dots$ and the received feedback $z_1, z_2, \dots \in \{0, 1\}$ in the following inductive way on $i \in \mathbb{N}$

$$\begin{aligned} x_1 &:= \alpha_1(w_1), & z_1 &:= \varphi(x_1, v_1), \\ x_{i+1} &:= \alpha_{i+1}(w_1, \dots, w_{i+1}, z_1, \dots, z_i), & z_{i+1} &:= \varphi(x_{i+1}, v_{i+1}). \end{aligned}$$

For each $a \in [0, 1]$, fix a binary representation $0.a_1a_2a_3\dots$ and define $\xi(a) := 0.a_1a_3a_5\dots$ and $\zeta(a) := 0.a_2a_4a_6\dots$. Notice that $\xi, \zeta: [0, 1] \rightarrow [0, 1]$ are independent with respect to the Lebesgue measure on $[0, 1]$ and that their (common) distribution is a uniform on $[0, 1]$. For each $x \in [0, 1]$, define $\psi_x: [0, 1] \rightarrow \{0, 1\}, u \mapsto \mathbf{1}_{[0, 1/4]}(x) + \mathbf{1}_{(1/4, 1/2]}(x) \cdot \mathbf{1}_{[0, 1-a]}(u) + \mathbf{1}_{(3/4, 1]}(x) \cdot \mathbf{1}_{[0, 1-a-2b]}(u)$. Define by induction on $i \in \mathbb{N}$ the following process

$$\begin{aligned} \tilde{x}_1 &:= \alpha_1(\zeta(w_1)), \\ \tilde{z}_1 &:= \varphi(\tilde{x}_1, \psi_{\tilde{x}_1}(\xi(w_1))), \\ \tilde{x}_{i+1} &:= \alpha_{i+1}(\zeta(w_1), \dots, \zeta(w_{i+1}), \tilde{z}_1, \dots, \tilde{z}_i), \\ \tilde{z}_{i+1} &:= \begin{cases} \varphi(\tilde{x}_{i+1}, v_{i+1}), & \tilde{x}_{i+1} \in (1/2, 3/4] \\ \varphi(\tilde{x}_{i+1}, \psi_{\tilde{x}_{i+1}}(\xi(w_{i+1}))), & \text{otherwise.} \end{cases} \end{aligned}$$

Since, for each $\epsilon \in [-1, 1]$ and each $i \in \mathbb{N}$,

$$\begin{aligned} P^\epsilon(z_i = 1 \mid x_i) &= \begin{cases} 1 & x_i \in [0, \frac{1}{4}] \\ 1 - a & x_i \in (\frac{1}{4}, \frac{1}{2}] \\ 1 - a - (1 + \epsilon) \cdot b & x_i \in (\frac{1}{2}, \frac{3}{4}] \\ 1 - a - 2 \cdot b & x_i \in (\frac{3}{4}, 1] \end{cases}, \\ P^\epsilon(\tilde{z}_i = 1 \mid \tilde{x}_i) &= \begin{cases} 1 & \tilde{x}_i \in [0, \frac{1}{4}] \\ 1 - a & \tilde{x}_i \in (\frac{1}{4}, \frac{1}{2}] \\ 1 - a - (1 + \epsilon) \cdot b & \tilde{x}_i \in (\frac{1}{2}, \frac{3}{4}] \\ 1 - a - 2 \cdot b & \tilde{x}_i \in (\frac{3}{4}, 1] \end{cases} \end{aligned}$$

it follows that, for each $\epsilon \in [-1, 1]$ and each $i \in \mathbb{N}$, the random variable $\tilde{x}_i$ has the same distribution as the random choice x_i made by the randomized algorithm α at time i when the underlying distribution is P^ϵ , i.e.,

$$P_{\tilde{x}_i}^\epsilon = P_{x_i}^\epsilon . \quad (50)$$

As with the process $x_1, x_2, \dots$, we have to count the number of times the process $\tilde{x}_1, \tilde{x}_2, \dots$ lands in the regions $(1/2, 3/4]$, $[0, 1/2]$ and $(3/4, 1]$ up to the time T . This can be done relying on the following random variables

$$\tilde{n}_1 := \sum_{i=1}^T \mathbf{1}_{(1/2, 3/4]}(\tilde{x}_i) , \quad \tilde{n}_2 := \sum_{i=1}^T \mathbf{1}_{[0, 1/2]}(\tilde{x}_i) , \quad \tilde{n}_3 := \sum_{i=1}^T \mathbf{1}_{(3/4, 1]}(\tilde{x}_i) .$$

Since, for each $\epsilon \in [-1, 1]$ and each $j \in \{1, 2, 3\}$,

$$E^\epsilon(\tilde{n}_j) = \sum_{i=1}^T P_{x_i}^\epsilon((1/2, 3/4]) \stackrel{(50)}{=} \sum_{i=1}^T P_{\tilde{x}_i}^\epsilon((1/2, 3/4]) = E^\epsilon(n_j) ,$$

to prove the claim (43), it is enough to prove that, for each $\epsilon \in [-1, 1]$,

$$E^{-\epsilon}(\tilde{n}_3) \geq E^\epsilon(\tilde{n}_3) - c_3 \cdot \epsilon \cdot T \cdot \sqrt{E^\epsilon(\tilde{n}_1)} .$$

We first prove the result when the sequence of randomization seeds is fixed, i.e., we suppose first that $\bar{w}_1, \bar{w}_2, \dots$ are such that $w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots$. For each $\epsilon \in [-1, 1]$, we consider the associated probability distribution Q^ϵ , defined as the conditional probability distribution $P^\epsilon(\cdot \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots)$. For each $t \in \mathbb{N}$, let $I_t := \{i \in \{1, \dots, t\} \mid \tilde{x}_i \in (1/2, 3/4]\}$, and for each $s \in \{1, \dots, t\}$, let

$$Z_{t,s} := \begin{cases} \emptyset & \text{if } s \notin I_t , \\ \mathbf{1}_{(1/2 < v_s)} & \text{if } s \in I_t . \end{cases}$$

Notice that for each $t_1, t_2 \in \mathbb{N}$ and each $s \in \{1, \dots, \min(t_1, t_2)\}$, we have that $Z_{t_1,s} = Z_{t_2,s}$. Then, for each $s \in \mathbb{N}$, it is well defined the random variable $Z_s := Z_{t,s}$, where $t \in \mathbb{N}$ is any number $t \geq s$. Define, for each $t \in \mathbb{N}$, the random vector $\bar{Z}_t := (Z_1, \dots, Z_t)$. Notice that, given that the sequence of randomization seeds is fixed and that, for each $s \in \mathbb{N}$, we have that $v_s \in \{1/4, 1/2, 3/4, 1\}$ (hence, for each $x \in (1/2, 3/4]$, it holds that $\mathbf{1}_{(1/2 < v_s)} = \mathbf{1}(x = v_s)$), the random vector $(\tilde{x}_1, \dots, \tilde{x}_T)$ is measurable with respect to the σ -algebra generated by $\bar{Z}_{T-1}$. Hence, for each $\epsilon \in [0, 1]$ and each $i \in \{1, \dots, T\}$, we

can deduce from Pinsker's inequality (see, e.g., (Tsybakov, 2008, Lemma 2.5)) that

$$Q^\epsilon(\tilde{x}_i \in (3/4, 1]) \leq Q^{-\epsilon}(\tilde{x}_i \in (3/4, 1]) + \sqrt{\frac{1}{2} \mathcal{D}_{\text{KL}}(Q_{\tilde{Z}_{T-1}}^\epsilon \parallel Q_{\tilde{Z}_{T-1}}^{-\epsilon})}, \quad (51)$$

where $\mathcal{D}_{\text{KL}}$ is the Kullback-Leibler divergence. Now, for each $t \in \mathbb{N}$ and each $\epsilon \in [0, 1]$, by the chain rule for Kullback-Leibler divergence (see, e.g., (Cover and Thomas, 2006, Theorem 2.5.3)), we have

$$\begin{aligned} \mathcal{D}_{\text{KL}}(Q_{\tilde{Z}_{t+1}}^\epsilon \parallel Q_{\tilde{Z}_{t+1}}^{-\epsilon}) &= \mathcal{D}_{\text{KL}}(Q_{(\tilde{Z}_t, Z_{t+1})}^\epsilon \parallel Q_{(\tilde{Z}_t, Z_{t+1})}^{-\epsilon}) \\ &= \mathcal{D}_{\text{KL}}(Q_{\tilde{Z}_t}^\epsilon \parallel Q_{\tilde{Z}_t}^{-\epsilon}) + \sum_{(\bar{z}, z) \in \{\emptyset, 0, 1\}^t \times \{\emptyset, 0, 1\}} \log \left(\frac{Q^\epsilon(Z_{t+1} = z \mid \bar{Z}_t = \bar{z})}{Q^{-\epsilon}(Z_{t+1} = z \mid \bar{Z}_t = \bar{z})} \right) \cdot Q^\epsilon(\bar{Z}_t = \bar{z} \cap Z_{t+1} = z). \end{aligned} \quad (52)$$

Notice that, for each $t \in \mathbb{N}$ and each $\epsilon \in [0, 1]$ we have

$$\begin{aligned} &\sum_{(\bar{z}, z) \in \{\emptyset, 0, 1\}^t \times \{\emptyset, 0, 1\}} \log \left(\frac{Q^\epsilon(Z_{t+1} = z \mid \bar{Z}_t = \bar{z})}{Q^{-\epsilon}(Z_{t+1} = z \mid \bar{Z}_t = \bar{z})} \right) \cdot Q^\epsilon(\bar{Z}_t = \bar{z} \cap Z_{t+1} = z) \\ &= \sum_{\substack{(\bar{z}, z) \in \{\emptyset, 0, 1\}^t \times \{\emptyset, 0, 1\} \\ t+1 \in I_{t+1}}} \log \left(\frac{Q^\epsilon(Z_{t+1} = z \mid \bar{Z}_t = \bar{z})}{Q^{-\epsilon}(Z_{t+1} = z \mid \bar{Z}_t = \bar{z})} \right) \cdot Q^\epsilon(\bar{Z}_t = \bar{z} \cap Z_{t+1} = z) \\ &= \left(\sum_{\substack{\bar{z} \in \{\emptyset, 0, 1\}^t \\ t+1 \in I_{t+1}}} Q^\epsilon(\bar{Z}_t = \bar{z}) \right) \cdot \sum_{z \in \{0, 1\}} \log \left(\frac{Q^\epsilon(\mathbf{1}(1/2 < v_{t+1}) = z)}{Q^{-\epsilon}(\mathbf{1}(1/2 < v_{t+1}) = z)} \right) \cdot Q^\epsilon(\mathbf{1}(1/2 < v_{t+1}) = z) \\ &= Q^\epsilon(\tilde{x}_{t+1} \in (1/2, 3/4]) \cdot \sum_{z \in \{0, 1\}} \log \left(\frac{Q^\epsilon(\mathbf{1}(1/2 < v_{t+1}) = z)}{Q^{-\epsilon}(\mathbf{1}(1/2 < v_{t+1}) = z)} \right) \cdot Q^\epsilon(\mathbf{1}(1/2 < v_{t+1}) = z). \end{aligned} \quad (53)$$

Algebraic calculations show that, for each $t \in \mathbb{N}$ and each $\epsilon \in [0, 1]$,

$$\begin{aligned}
& \sum_{z \in \{0,1\}} \log \left(\frac{Q^\epsilon(\mathbf{1}(1/2 < v_{t+1}) = z)}{Q^{-\epsilon}(\mathbf{1}(1/2 < v_{t+1}) = z)} \right) \cdot Q^\epsilon(\mathbf{1}(1/2 < v_{t+1}) = z) \\
&= \log \left(\frac{Q^\epsilon(\frac{1}{2} < v_{t+1})}{Q^{-\epsilon}(\frac{1}{2} < v_{t+1})} \right) \cdot Q^\epsilon\left(\frac{1}{2} < v_{t+1}\right) + \log \left(\frac{Q^\epsilon(\frac{1}{2} \geq v_{t+1})}{Q^{-\epsilon}(\frac{1}{2} \geq v_{t+1})} \right) \cdot Q^\epsilon\left(\frac{1}{2} \geq v_{t+1}\right) \\
&= \log \left(\frac{1-a-(1+\epsilon) \cdot b}{1-a-(1-\epsilon) \cdot b} \right) \cdot (1-a-(1+\epsilon) \cdot b) + \log \left(\frac{a+(1+\epsilon) \cdot b}{a+(1-\epsilon) \cdot b} \right) \cdot (a+(1+\epsilon) \cdot b) \\
&\leq \frac{4 \cdot b^2 \cdot \epsilon^2}{(1-a-(1-\epsilon) \cdot b) \cdot (a+(1-\epsilon) \cdot b)} \leq \frac{4 \cdot b^2 \cdot \epsilon^2}{a \cdot (1-a-2b)} = 2 \cdot c_3^2 \cdot \epsilon^2.
\end{aligned} \tag{54}$$

Putting (52), (53) and (54) together, we obtain that, for each $t \in \mathbb{N}$ and each $\epsilon \in [0, 1]$,

$$\mathcal{D}_{\text{KL}}(Q_{\bar{Z}_{t+1}}^\epsilon \parallel Q_{\bar{Z}_{t+1}}^{-\epsilon}) \leq \mathcal{D}_{\text{KL}}(Q_{Z_1}^\epsilon \parallel Q_{Z_1}^{-\epsilon}) + 2 \cdot c_3^2 \cdot \epsilon^2 \cdot \sum_{s=1}^t Q^\epsilon(\tilde{x}_{s+1} \in (1/2, 3/4]) . \tag{55}$$

With the same technique used above, for each $\epsilon \in [0, 1]$, we can prove that

$$\mathcal{D}_{\text{KL}}(Q_{Z_1}^\epsilon \parallel Q_{Z_1}^{-\epsilon}) \leq 2 \cdot c_3^2 \cdot \epsilon^2 \cdot Q^\epsilon(\tilde{x}_1 \in (1/2, 3/4]) . \tag{56}$$

For each $t \in \{1, \dots, T\}$, putting (55) and (56) together, we obtain

$$\begin{aligned}
\mathcal{D}_{\text{KL}}(Q_{\bar{Z}_t}^\epsilon \parallel Q_{\bar{Z}_t}^{-\epsilon}) &\stackrel{(55)+(56)}{\leq} 2 \cdot c_3^2 \cdot \epsilon^2 \cdot \sum_{s=1}^t Q^\epsilon(\tilde{x}_s \in (1/2, 3/4]) \\
&\leq 2 \cdot c_3^2 \cdot \epsilon^2 \cdot E^\epsilon(\tilde{n}_1 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) .
\end{aligned} \tag{57}$$

Now, (51) and (57) imply that, for each $\epsilon \in [0, 1]$ and each $i \in \{1, \dots, T\}$,

$$Q^\epsilon(\tilde{x}_i \in (3/4, 1]) \leq Q^{-\epsilon}(\tilde{x}_i \in (3/4, 1]) + c_3 \cdot \epsilon \cdot \sqrt{E^\epsilon(\tilde{n}_1 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots)} . \tag{58}$$

Taking the sum of (58) over $i \in \{1, \dots, T\}$, we obtain that for each $\epsilon \in [0, 1]$,

$$\begin{aligned} E^{-\epsilon}(\tilde{n}_3 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) \\ \geq E^\epsilon(\tilde{n}_3 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) - c_3 \cdot \epsilon \cdot T \cdot \sqrt{E^\epsilon(\tilde{n}_1 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots)} . \end{aligned} \quad (59)$$

Now, since the sequence $\bar{w}_1, \bar{w}_2, \dots$ of randomization seeds has been arbitrarily chosen, for each $\epsilon \in [0, 1]$, using the fact that $P_{(w_t)_{t \in \mathbb{N}}}^\epsilon = P_{(w_t)_{t \in \mathbb{N}}}^{-\epsilon}$ and Jensen's inequality, we have that

$$\begin{aligned} E^{-\epsilon}(\tilde{n}_3) &= \int_{[0,1]^\mathbb{N}} E^{-\epsilon}(\tilde{n}_3 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) dP_{(w_t)_{t \in \mathbb{N}}}^{-\epsilon}(\bar{w}_1, \bar{w}_2, \dots) \\ &\stackrel{(49)}{=} \int_{[0,1]^\mathbb{N}} E^{-\epsilon}(\tilde{n}_3 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) dP_{(w_t)_{t \in \mathbb{N}}}^\epsilon(\bar{w}_1, \bar{w}_2, \dots) \\ &\stackrel{(59)}{\geq} \int_{[0,1]^\mathbb{N}} E^\epsilon(\tilde{n}_3 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) dP_{(w_t)_{t \in \mathbb{N}}}^\epsilon(\bar{w}_1, \bar{w}_2, \dots) \\ &\quad - c_3 \cdot \epsilon \cdot T \cdot \int_{[0,1]^\mathbb{N}} \sqrt{E^\epsilon(\tilde{n}_1 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots)} dP_{(w_t)_{t \in \mathbb{N}}}^\epsilon(\bar{w}_1, \bar{w}_2, \dots) \\ (\text{by Jensen}) \quad &\geq \int_{[0,1]^\mathbb{N}} E^\epsilon(\tilde{n}_3 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) dP_{(w_t)_{t \in \mathbb{N}}}^\epsilon(\bar{w}_1, \bar{w}_2, \dots) \\ &\quad - c_3 \cdot \epsilon \cdot T \cdot \sqrt{\int_{[0,1]^\mathbb{N}} E^\epsilon(\tilde{n}_1 \mid w_1 = \bar{w}_1, w_2 = \bar{w}_2, \dots) dP_{(w_t)_{t \in \mathbb{N}}}^\epsilon(\bar{w}_1, \bar{w}_2, \dots)} \\ &= E^\epsilon(\tilde{n}_3) - c_3 \cdot \epsilon \cdot \sqrt{E^\epsilon(\tilde{n}_1)} . \end{aligned}$$

□

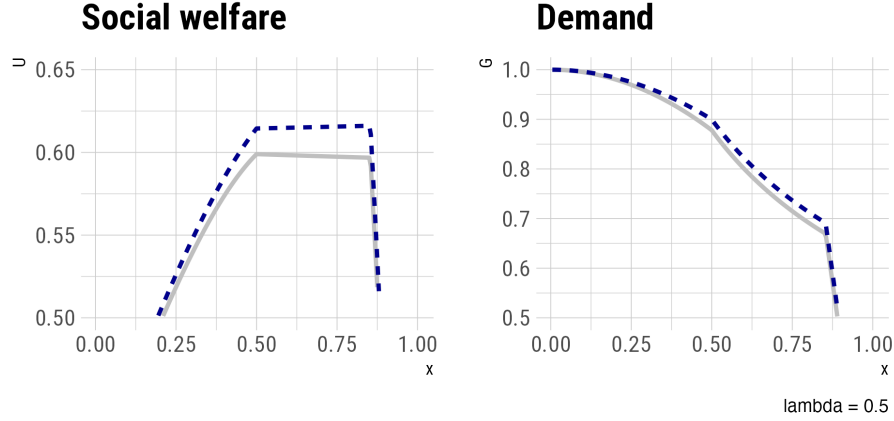
B.2 Theorem 3 (Lower bound on regret for the concave case)

Defining a family of distributions for v Define $\bar{h} := \frac{1-\sqrt{1-\lambda}}{2}$ and notice that $0 < \bar{h} < \frac{1}{2}$. Define $\bar{\eta} := (\bar{h} \cdot (1 - \bar{h})^{1-\lambda} \cdot (1 - \lambda))^{-1}$ and $\bar{\epsilon} := \frac{1}{2} \cdot \min(\bar{\eta}, \frac{2}{3} \cdot 2^{-\lambda})$. For each $\epsilon \in (-\bar{\epsilon}, \bar{\epsilon})$ and each $x \in [0, 1]$, define

$$f^\epsilon(x) := \bar{c} \cdot \left((2^{2-\lambda} - 8 \cdot \bar{h} \cdot \epsilon) \cdot x \cdot \mathbf{1}_{[0, \frac{1}{2})}(x) + \frac{1}{x^{2-\lambda}} \cdot \mathbf{1}_{[\frac{1}{2}, 1-\bar{h}]}(x) + (\bar{\eta} + \epsilon) \cdot \mathbf{1}_{(1-\bar{h}, 1]}(x) \right) ,$$

where $\bar{c}$ is such that $\int_0^1 f^0(x) dx = 1$. For each $\epsilon \in (-\bar{\epsilon}, \bar{\epsilon})$, note that f^ϵ is a density function on $[0, 1]$, i.e., a non-negative function whose integral is 1. For each $\epsilon \in (-\bar{\epsilon}, \bar{\epsilon})$, let μ^ϵ be the probability measure whose density is f^ϵ , and define $\mathbf{G}^\epsilon$ and $\mathbf{U}^\epsilon$ as the

Figure 3: Construction for proving the lower bound on regret for the concave case



demand function and the expected social welfare associated to μ^ϵ , respectively (see Figure 3 for an illustration).

Properties of U Define also $\bar{x} := \frac{1}{2} \cdot (\frac{1}{2} + (1 - \bar{h})) = \frac{3}{4} - \frac{\bar{h}}{2}$ and $\bar{m} := \frac{1 - \sqrt{1 - \lambda}}{8} \cdot (1 - \lambda)^{3/2}$. Notice that, for all $\epsilon \in (-\bar{\epsilon}, \bar{\epsilon})$, we have:

- U^ϵ is continuous and concave.
- U^ϵ is strictly increasing in $[0, \frac{1}{2}]$, linear in $[\frac{1}{2}, 1 - \bar{h}]$ with slope $(1 - \lambda) \cdot \bar{h} \cdot \epsilon$, and strictly decreasing on $[1 - \bar{h}, 1]$, which in particular implies that the maximum of U^ϵ is at $1 - \bar{h}$ if $\epsilon > 0$, and at $\frac{1}{2}$ if $\epsilon < 0$.
- If $\epsilon > 0$, then $U^\epsilon(1 - \bar{h}) - \max_{x \in [0, \bar{x}]} U^\epsilon(x) = \bar{m} \cdot |\epsilon| = U^{-\epsilon}(\frac{1}{2}) - \max_{x \in [\bar{x}, 1]} U^{-\epsilon}(x)$.

Now, consider the sequence of individual valuations $v_1, v_2, \dots \in [0, 1]$, and assume that, for each $\epsilon \in (-\bar{\epsilon}, \bar{\epsilon})$, when the underlying distribution is P^ϵ , this sequence is i.i.d. (independent of the randomization used by the algorithm) with common distribution μ^ϵ . The previous list of properties implies that, for each $\epsilon \in (0, \bar{\epsilon})$ (resp., $\epsilon \in (-\bar{\epsilon}, 0)$), when the underlying distribution is P^ϵ , the expected instantaneous regret at time t is at least $\bar{m} \cdot |\epsilon|$ if the learner plays in the region $\bar{I} := [0, \bar{x}]$ (resp., in the region $\bar{J} := [\bar{x}, 1]$). It follows that, in order not to suffer linear regret, the learner has to discriminate the sign of ϵ .

Intuition for the proof Now, the high-level idea is that in order to discriminate the sign of ϵ , the learner needs on the order of $\frac{1}{\epsilon^2}$ observations. Therefore, for a number of periods on the order of $\frac{1}{\epsilon^2}$, the algorithm is playing “in the dark,” and thus suffers an expected regret on the order of $\epsilon \cdot T$, or, equivalently, on the order of $\sqrt{T}$, when the underlying distribution is between P^ϵ or $P^{-\epsilon}$.

Defining constants We now formalize this idea. Let

$$\gamma := \left(\int_0^{1/2} \left(\frac{16\bar{h}}{2^{2-\lambda} - 8\bar{h}\bar{\epsilon}} \right)^2 f^{-\bar{\epsilon}}(x) dx + \int_{1-\bar{h}}^1 \left(\frac{2}{\bar{\eta} - \bar{\epsilon}} \right)^2 f^{\bar{\epsilon}}(x) dx \right)^{1/2} > 0$$

Let $\bar{M} > 0$ such that $2 \cdot \sqrt{\frac{\sqrt{2}}{3} \cdot \frac{\gamma \bar{M}}{\bar{m}}} = 1$. Let $M \in (0, \bar{M})$ such that

$$k := \sqrt{\frac{\frac{M}{\bar{m}}}{\frac{\sqrt{2}}{3} \cdot \gamma}} < \bar{\epsilon}.$$

From now on, fix a time horizon $T \in \mathbb{N}$ and let $\epsilon := \frac{k}{\sqrt{T}}$. In the following we use the notation E^ϵ (resp., $E^{-\epsilon}$) to denote the expectation with respect to the probability measure P^ϵ (resp., $P^{-\epsilon}$). Let $x_1, x_2, \dots$ be the policies chosen by the algorithm. Note that, since the algorithm bases its decision at time t only on the (partial) knowledge of $v_1, \dots, v_{t-1}$ and some independent randomization, there exists a (measurable) function $\varphi_t: [0, 1]^{t-1} \rightarrow [0, 1]$ such that

$$E^\epsilon(\mathbf{1}(x_t \in \bar{I}) \mid v_1, \dots, v_{t-1}) = \varphi_t(v_1, \dots, v_{t-1}) = E^{-\epsilon}(\mathbf{1}(x_t \in \bar{I}) \mid v_1, \dots, v_{t-1}).$$

Then, for each time t , it holds

$$\begin{aligned} |E^\epsilon(\mathbf{1}(x_t \in \bar{I})) - E^{-\epsilon}(\mathbf{1}(x_t \in \bar{I}))| &= |E^\epsilon(\varphi_t(v_1, \dots, v_{t-1})) - E^{-\epsilon}(\varphi_t(v_1, \dots, v_{t-1}))| \\ &\leq \left\| \bigotimes_{s=1}^{t-1} \mu^\epsilon - \bigotimes_{s=1}^{t-1} \mu^{-\epsilon} \right\|_{\text{TV}} = (\star) \end{aligned}$$

Relating choice probabilities for positive and negative ϵ By Pinsker’s inequality and the fact that the Kullback-Leibler divergence is upper bounded by the χ^2 -

divergence, it follows that

$$\begin{aligned}
 (\star) &\leq \sqrt{\frac{\mathcal{D}_{\text{KL}}(\otimes_{s=1}^{t-1} \mu^{-\epsilon}, \otimes_{s=1}^{t-1} \mu^{\epsilon})}{2}} = \sqrt{\frac{(t-1) \cdot \mathcal{D}_{\text{KL}}(\mu^{-\epsilon}, \mu^{\epsilon})}{2}} \leq \sqrt{\frac{(t-1) \cdot \mathcal{D}_{\chi^2}(\mu^{-\epsilon}, \mu^{\epsilon})}{2}} \\
 &= \sqrt{\frac{t-1}{2} \int_0^1 \left| \frac{f^{\epsilon}(x)}{f^{-\epsilon}(x)} - 1 \right|^2 f^{-\epsilon}(x) dx} = (\star\star)
 \end{aligned}$$

Now, noticing that

$$\begin{aligned}
 \int_0^1 \left| \frac{f^{\epsilon}(x)}{f^{-\epsilon}(x)} - 1 \right|^2 f^{-\epsilon}(x) dx &= \int_0^{1/2} \left(\frac{16\bar{h}\epsilon}{2^{2-\lambda} + 8\bar{h}\epsilon} \right)^2 f^{-\epsilon}(x) dx + \int_{1-\bar{h}}^1 \left(\frac{2\epsilon}{\bar{\eta} - \epsilon} \right)^2 f^{-\epsilon}(x) dx \\
 &\leq \left(\int_0^{1/2} \left(\frac{16\bar{h}}{2^{2-\lambda} - 8\bar{h}\bar{\epsilon}} \right)^2 f^{-\bar{\epsilon}}(x) dx + \int_{1-\bar{h}}^1 \left(\frac{2}{\bar{\eta} - \bar{\epsilon}} \right)^2 f^{\bar{\epsilon}}(x) dx \right) \cdot \epsilon^2 \\
 &= \gamma^2 \cdot \epsilon^2,
 \end{aligned}$$

it follows that

$$(\star\star) \leq \gamma \cdot \epsilon \cdot \sqrt{\frac{t-1}{2}}.$$

Summing over $t = 1, 2, \dots, T$, we obtain

$$\left| E^{\epsilon} \left(\sum_{t=1}^T \mathbf{1}(x_t \in \bar{I}) \right) - E^{-\epsilon} \left(\sum_{t=1}^T \mathbf{1}(x_t \in \bar{I}) \right) \right| \leq \frac{\sqrt{2}}{3} \cdot \gamma \cdot \epsilon \cdot T^{3/2} = \frac{\sqrt{2}}{3} \cdot \gamma \cdot k \cdot T.$$

Upper bound on regret for $\epsilon > 0$ implies lower bound on regret for $-\epsilon$.

Now, suppose that in the scenario determined by P^{ϵ} the algorithm suffer a regret $R_T^{\epsilon} \leq M \cdot \sqrt{T}$. Then

$$M \cdot \sqrt{T} \geq R_T^{\epsilon} \geq \bar{m} \cdot \epsilon \cdot \sum_{t=1}^T E^{\epsilon}(\mathbf{1}(x_t \in \bar{I})) = \bar{m} \cdot \frac{k}{\sqrt{T}} \cdot \sum_{t=1}^T E^{\epsilon}(\mathbf{1}(x_t \in \bar{I})).$$

and rearranging

$$\sum_{t=1}^T E^{\epsilon}(\mathbf{1}(x_t \in \bar{I})) \leq \frac{M \cdot T}{\bar{m} \cdot k}$$

It follows that the expected number of times the algorithm plays in the (correct) region I when the underlying scenario is determined by $P^{-\epsilon}$ is

$$\begin{aligned} \sum_{t=1}^T E^{-\epsilon}(\mathbf{1}(x_t \in \bar{I})) &= \left(\sum_{t=1}^T E^{-\epsilon}(\mathbf{1}(x_t \in \bar{I})) - \sum_{t=1}^T E^{\epsilon}(\mathbf{1}(x_t \in \bar{I})) \right) + \sum_{t=1}^T E^{\epsilon}(\mathbf{1}(x_t \in \bar{I})) \\ &\leq \left(\frac{\sqrt{2}}{3} \cdot \gamma \cdot k + \frac{M}{\bar{m} \cdot k} \right) \cdot T \end{aligned}$$

The last inequality implies that the expected number of times that the algorithm plays in the (wrong) region $J = I^c$ when the underlying scenario is determined by $P^{-\epsilon}$ is lower bounded by

$$\sum_{t=1}^T E^{-\epsilon}(\mathbf{1}(x_t \in \bar{J})) = \sum_{t=1}^T E^{-\epsilon}(\mathbf{1}(x_t \notin \bar{I})) \geq \left(1 - \left(\frac{\sqrt{2}}{3} \cdot \gamma \cdot k + \frac{M}{\bar{m} \cdot k} \right) \right) \cdot T,$$

which implies that the regret the algorithm suffers in the scenario determined by $P^{-\epsilon}$ is lower bounded by

$$\begin{aligned} R_T^{-\epsilon} &\geq \bar{m} \cdot \epsilon \cdot \sum_{t=1}^T E^{-\epsilon}(\mathbf{1}(x_t \in J)) = \bar{m} \cdot \frac{k}{\sqrt{T}} \cdot \sum_{t=1}^T E^{-\epsilon}(\mathbf{1}(x_t \in J)) \\ &\geq \bar{m} \cdot \frac{k}{\sqrt{T}} \cdot \left(1 - \left(\frac{\sqrt{2}}{3} \cdot \gamma \cdot k + \frac{M}{\bar{m} \cdot k} \right) \right) \cdot T = \bar{m} \cdot k \cdot \left(1 - 2\sqrt{\frac{\sqrt{2} \cdot \gamma \cdot M}{3 \cdot \bar{m}}} \right) \cdot \sqrt{T}. \end{aligned}$$

Putting everything together, any algorithm suffers at least $\min \left(M, \bar{m} \cdot k \cdot \left(1 - 2 \cdot \sqrt{\frac{\sqrt{2} \cdot \gamma \cdot M}{3 \cdot \bar{m}}} \right) \right) \cdot \sqrt{T}$ regret, in at least one scenario between the ones determined by P^{ϵ} and $P^{-\epsilon}$. Recalling that our choice of M implies $1 - 2\sqrt{\frac{\sqrt{2} \cdot \gamma \cdot M}{3 \cdot \bar{m}}} > 0$, the conclusion follows.

B.3 Theorem 4 (Stochastic upper bound on regret of Dyadic Search for Social Welfare)

For the sake of simplicity, we assume that $\mathbf{U}$ admits a unique maximizer $x^* \in [0, 1]$ (the other cases can be treated similarly and, actually, they ended up having better constants in the final regret guarantees).

For each epoch $\tau = 1, 2, \dots$, we refer to the three current l (left), c (center) and r (right) points of the corresponding epoch τ using l_{τ}, c_{τ} and r_{τ} , respectively. For any

time t , the epoch to which the time t belongs is denoted τ_t . The length of an interval J is denoted $|J|$, while the number of elements in a finite set A is denoted $\#A$.

Consider a family $(v_{x,i})_{x \in [0,1], i \in \mathbb{N}}$ of random variables such that, for each $x \in [0,1]$, the sequence $(v_{x,i})_{i \in \mathbb{N}}$ is i.i.d. with the same distribution as $(v_i)_{i \in \mathbb{N}}$. With these random variables, we can define the auxiliary family $(y_{x,i})_{x \in [0,1], i \in \mathbb{N}} := (\mathbf{1}(x \leq v_{x,i}))_{x \in [0,1], i \in \mathbb{N}}$. We assume that, whenever we select a policy $x \in [0,1]$ at time t , we observe $\mathbf{1}(x \leq v_{x,n_t(x)})$ (recall that $n_t(x) = \sum_{s=1}^t \mathbf{1}(x_s = x)$) instead of $\mathbf{1}(x \leq v_t)$. This does not change anything in expectation, but will be useful in what follows.

The next lemma states that Algorithm 2 maintains confidence intervals containing the differences of the welfare function (among left, center and right points) with high probability.

Lemma 1 (Confidence intervals contain true welfare differences with high probability). *There exists a constant $\tilde{C} \in (0, 20]$ such that, for every time horizon T and any $\delta \in (0, 1)$, if the learner runs Algorithm 2 with confidence parameter δ , then the probability of the event*

$$\mathcal{E} := \bigcap_{t=1}^T \left(\{U(c_{\tau_t}) - U(l_{\tau_t}) \in J_t(l_{\tau_t}, c_{\tau_t})\} \cap \{U(r_{\tau_t}) - U(c_{\tau_t}) \in J_t(c_{\tau_t}, l_{\tau_t})\} \cap \{U(r_{\tau_t}) - U(l_{\tau_t}) \in J_t(r_{\tau_t}, l_{\tau_t})\} \right)$$

is lower bounded by $1 - \tilde{C} \cdot T^2 \cdot \delta$.

The proof of this lemma can be found in B.3.1.

The following lemma establishes the rate of shrinking of the length of the confidence intervals as the length of an epoch increases.

Lemma 2 (Confidence intervals shrink with epoch length). *For any $\delta \in (0, 1)$, if the learner runs Algorithm 2 with confidence parameter δ then, for any time t ,*

$$\max(|J_t(l_{\tau_t}, c_{\tau_t})|, |J_t(c_{\tau_t}, r_{\tau_t})|, |J_t(l_{\tau_t}, r_{\tau_t})|) \leq \frac{\tilde{c}_\delta}{\sqrt{t - t_{\tau_t-1}}}, \quad (60)$$

whenever $t - t_{\tau_t-1} \geq \tilde{n}$, where $\tilde{n} = 10$ and $\tilde{c}_\delta = 72 \cdot \sqrt{10} \cdot (\sqrt{2 \log(2/\delta)} + 4)$.

The proof of this lemma can be found in B.3.2.

Lemma 1 and Lemma 2 allow us to prove Theorem 4, which closely follows the proof given in (Bachoc et al., 2022b).

B.3.1 Lemma 1 (Confidence intervals contain true welfare differences with high probability)

Proof. For each $n \in \mathbb{N}$, let $\mathcal{D}_n := \{k \cdot 2^{-n} \mid k \in \mathbb{Z}\}$, let $\mathcal{D}_n^* := \{x_{n,1}, \dots, x_{n,10}\} \subset \mathcal{D}_n$ such that

$$x_{n,1} < \dots < x_{n,5} \leq x^* \leq x_{n,6} < \dots < x_{n,10}$$

and $x_{n,j+1} - x_{n,j} \leq 2^{-n}$, for all $j \in \{1, \dots, 9\}$. Define $\mathcal{D} := \bigcup_{n=1}^T \mathcal{D}_n^* \cap (0, 1)$. Consider the following events

$$\mathcal{E}' := \bigcap_{\substack{n, t \in \{1, \dots, T\} \\ j \in \{1, \dots, 10\}}} \left\{ \left| \frac{1}{t} \sum_{s=1}^t y_{x_{n,j}, s} - \mathbf{G}(x_{n,j}) \right| \leq \sqrt{\frac{1}{2t} \log \left(\frac{2}{\delta} \right)} \right\}$$

$$\mathcal{E}'' := \bigcap_{\substack{n \in \{1, \dots, T\} \\ m \in \{1, \dots, \lfloor \log_2(T) \rfloor\} \\ j \in \{1, \dots, 9\}}} \left\{ \left| \frac{1}{2^m} \sum_{i=1}^{2^m-1} y_{x_{n,j} + \frac{i}{2^{n+m}}, 1} - \frac{1}{x_{n,j+1} - x_{n,j}} \cdot \int_{x_{n,j}}^{x_{n,j+1}} \mathbf{G}(x) dx \right| \leq \sqrt{\frac{1}{2 \cdot 2^m} \log \left(\frac{2}{\delta} \right)} + \frac{2}{2^m} \right\}$$

and note that $\mathcal{E} \subset \mathcal{E}' \cap \mathcal{E}''$, since, in the event $\mathcal{E}' \cap \mathcal{E}''$, Algorithm 2 will query only points in $\mathcal{D}^*$, given that it uses only a subset of the estimates in the definition of $\mathcal{E}'$ and $\mathcal{E}''$ to build its own estimates (in particular, due to the ties breaking rules, to estimate the integral terms it will only use the *first* query of the relevant dyadic points). Now, notice that for each $n \in \{1, \dots, n\}$, each $m \in \{1, \dots, \lfloor \log_2(T) \rfloor\}$ and each $j \in \{1, \dots, 9\}$ we have

$$\begin{aligned} & \left\{ \left| \frac{1}{2^m} \sum_{i=1}^{2^m-1} y_{x_{n,j} + \frac{i}{2^{n+m}}, 1} - \frac{1}{x_{n,j+1} - x_{n,j}} \cdot \int_{x_{n,j}}^{x_{n,j+1}} \mathbf{G}(x) dx \right| > \sqrt{\frac{1}{2 \cdot 2^m} \log \left(\frac{2}{\delta} \right)} + \frac{2}{2^m} \right\} \\ & \subset \left\{ \left| \frac{1}{2^m} \sum_{i=1}^{2^m-1} y_{x_{n,j} + \frac{i}{2^{n+m}}, 1} - \frac{1}{2^m} \sum_{i=1}^{2^m-1} \mathbf{G} \left(x_{n,j} + \frac{i}{2^{n+m}} \right) \right| > \sqrt{\frac{1}{2 \cdot 2^m} \log \left(\frac{2}{\delta} \right)} \right\} \\ & \cup \left\{ \left| \frac{1}{2^m} \sum_{i=1}^{2^m-1} \mathbf{G} \left(x_{n,j} + \frac{i}{2^{n+m}} \right) - \frac{1}{x_{n,j+1} - x_{n,j}} \cdot \int_{x_{n,j}}^{x_{n,j+1}} \mathbf{G}(x) dx \right| > \frac{2}{2^m} \right\} \\ & = \left\{ \left| \frac{1}{2^m} \sum_{i=1}^{2^m-1} y_{x_{n,j} + \frac{i}{2^{n+m}}, 1} - \frac{1}{2^m} \sum_{i=1}^{2^m-1} \mathbf{G} \left(x_{n,j} + \frac{i}{2^{n+m}} \right) \right| > \sqrt{\frac{1}{2 \cdot 2^m} \log \left(\frac{2}{\delta} \right)} \right\} \end{aligned}$$

Algorithm 5 Index selection

```

1: for  $s = 1, 2, \dots$  do
2:   Let  $k_s = \min \left( \operatorname{argmax}_{k \in [5]} f_k(m_k(s-1)) \right)$ 
3:    $m_{k_s}(s) = m_{k_s}(s-1) + 1$ 
4:   for  $i \in [5] \setminus \{k_s\}$  do
5:      $m_i(s) = m_i(s-1)$ 
6:   end for
7: end for

```

where the last equality follows from

$$\begin{aligned}
& \left| \frac{1}{2^m} \sum_{i=1}^{2^m-1} \mathbf{G} \left(x_{n,j} + \frac{i}{2^{n+m}} \right) - \frac{1}{x_{n,j+1} - x_{n,j}} \cdot \int_{x_{n,j}}^{x_{n,j+1}} \mathbf{G}(x) dx \right| \\
& \leq \sum_{i=1}^{2^m-1} \int_{x_{n,j} + \frac{i-1}{2^{n+m}}}^{x_{n,j} + \frac{i}{2^{n+m}}} \left(\mathbf{G}(x) - \mathbf{G} \left(x_{n,j} + \frac{i}{2^{n+m}} \right) \right) dx + \frac{1}{2^m} \\
& \leq \sum_{i=1}^{2^m-1} \int_{x_{n,j} + \frac{i-1}{2^{n+m}}}^{x_{n,j} + \frac{i}{2^{n+m}}} \left(\mathbf{G} \left(x_{n,j} + \frac{i-1}{2^{n+m}} \right) - \mathbf{G} \left(x_{n,j} + \frac{i}{2^{n+m}} \right) \right) dx + \frac{1}{2^m} \\
& \leq \frac{1}{2^m} \cdot (\mathbf{G}(x_{n,j}) - \mathbf{G}(x_{n,j+1})) + \frac{1}{2^m} \leq \frac{2}{2^m}
\end{aligned}$$

By De Morgan's laws, a union bound and Hoeffding's inequality, we have $P(\mathcal{E}^c) \leq P((\mathcal{E}')^c) + P((\mathcal{E}'')^c) \leq 20 \cdot T^2 \cdot \delta$. $\square$

B.3.2 Lemma 2 (Confidence intervals shrink with epoch length)

We break the proof of Lemma 2 in several steps. Let $d_1, d_2, d_3, d_4, d_5 > 0$ be constants. For each $k \in \{1, 2, 3\}$, define

$$f_k : \{0, 1, 2, \dots\} \rightarrow [0, +\infty], \quad n \mapsto \frac{d_k}{\sqrt{n}}$$

and for each $k \in \{4, 5\}$ define

$$f_k : \{0, 1, 2, \dots\} \rightarrow [0, +\infty], \quad n \mapsto \frac{d_4}{\sqrt{2^{\lfloor \log_2(n+1) \rfloor} - 1}} + \frac{d_5}{2^{\lfloor \log_2(n+1) \rfloor}},$$

with the usual convention that $a/0 = +\infty$, for any $a > 0$. Suppose that $m_1(0), m_2(0), m_3(0), m_4(0), m_5(0) \in \{0, 1, 2, \dots\}$ and consider Algorithm 5.

The following lemma holds.

Lemma 3. Consider Algorithm [5](#) and the notation defined therein. For each $s \in \mathbb{N}$ there exists an index $i \in [5]$ for which $m_i(s) \geq \lceil s/5 \rceil$.

Proof. Let $s \in \mathbb{N}$ and suppose by contradiction that for each $k \in [5]$ it holds that $m_k(s) < s/5$. Then

$$s \leq \sum_{k=1}^5 m_k(s) \leq 5 \cdot \max_{k \in [5]} m_k(s) < 5 \cdot \frac{s}{5} = s ,$$

which is a contradiction. It follows that there exists $k \in [5]$ for which $m_k(s) \geq s/5$, which also implies $m_k(s) \geq \lceil s/5 \rceil$. Given that s was arbitrarily chosen, the conclusion follows. $\square$

Notice that, for each $n \in \{0, 1, 2, \dots\}$, we have

$$\frac{d_4}{\sqrt{n}} \leq \frac{d_4}{\sqrt{2^{\lfloor \log_2(n+1) \rfloor} - 1}} \leq \frac{2d_4}{\sqrt{n}}$$

and

$$0 \leq \frac{d_5}{\sqrt{2^{\lfloor \log_2(n+1) \rfloor}}} \leq \frac{2d_5}{n} ,$$

which implies that, for each $k \in [5]$ and each $n \in \{0, 1, 2, \dots\}$

$$\frac{d_k}{\sqrt{n}} \leq f_k(n) \leq \frac{D_k}{\sqrt{n}}$$

where $D_1 = d_1, D_2 = d_2, D_3 = d_3, D_4 = D_5 = 2(d_4 + d_5)$.

The following lemma holds.

Lemma 4. Consider Algorithm [5](#) and the notation defined therein. For any $i, j \in [5]$ and any $s \in \mathbb{N}$ it holds

$$m_i(s) \geq \left(\frac{d_i}{D_j} \right)^2 (m_j(s) - 1) .$$

Proof. Let $i, j \in [5]$. Suppose by contradiction that the conclusion does not hold. Then there exists a smallest $s \in \{0, 1, 2, \dots\}$ for which

$$m_i(s) < \left(\frac{d_i}{D_j} \right)^2 (m_j(s) - 1) ,$$

which we call s_0 . Notice that $s_0 \neq 0$. Then, the fact that

$$m_i(s_0 - 1) \geq \left(\frac{d_i}{D_j} \right)^2 (m_j(s_0 - 1) - 1),$$

implies that at time s_0 the algorithm selected $k_{s_0} = j$, which in turn implies that $m_i(s_0 - 1) = m_i(s_0)$ and $m_j(s_0 - 1) = m_j(s_0) - 1$. It follows that

$$\left(\frac{d_i}{D_j} \right)^2 m_j(s_0 - 1) = \left(\frac{d_i}{D_j} \right)^2 (m_j(s_0) - 1) > m_i(s_0) = m_i(s_0 - 1),$$

Rearranging, we get

$$m_j(s_0 - 1) > \left(\frac{D_j}{d_i} \right)^2 m_i(s_0 - 1).$$

from which it follows that

$$f_j(m_j(s_0 - 1)) \leq \frac{D_j}{\sqrt{m_j(s_0 - 1)}} < \frac{d_i}{\sqrt{m_i(s_0 - 1)}} \leq f_i(m_i(s_0 - 1)).$$

This last inequality implies that at time s_0 the algorithm should have chosen the index i and not the index j , which is a contradiction. $\square$

Combining the last two lemmas we can prove the following result.

Lemma 5. *Consider Algorithm 5 and the notation defined therein. Then, for any $s \geq 5$ it holds that*

$$\max_{k \in [5]} f_k(m_k(s)) \leq \frac{D}{\sqrt{s - 5}}$$

where $D = \sqrt{5} \cdot (\max_{j \in [5]} D_j) \cdot (\max_{k \in [5]} \frac{D_k}{d_k})$.

Proof. Let $s \geq 5$. Pick $j \in [5]$ such that $m_j(s) \geq \lceil s/5 \rceil$ (which does exist by Lemma 3). Then, by Lemma 4

$$\begin{aligned} \max_{k \in [5]} f_k(m_k(s)) &\leq \max_{k \in [5]} \frac{D_k}{\sqrt{m_k(s)}} \leq \max_{k \in [5]} \frac{D_k}{\sqrt{\left(\frac{d_k}{D_j} \right)^2 (m_j(s) - 1)}} \\ &= D_j \cdot \max_{k \in [5]} \left(\frac{D_k}{d_k} \right) \frac{1}{\sqrt{m_j(s) - 1}} \leq D_j \cdot \max_{k \in [5]} \left(\frac{D_k}{d_k} \right) \frac{1}{\sqrt{\lceil s/5 \rceil - 1}} \leq \frac{D}{\sqrt{s - 5}}. \end{aligned}$$

We are now ready for the proof of Lemma 2.

Proof of Lemma 2. It is enough to notice that Algorithm 2 with confidence parameter $\delta \in (0, 1)$ relies, inside each epoch, on the same routine given by Algorithm 5 with $d_1 = l \cdot \sqrt{\frac{\log(2/\delta)}{2}}$, $d_2 = c \cdot \sqrt{\frac{\log(2/\delta)}{2}}$, $d_3 = r \cdot \sqrt{\frac{\log(2/\delta)}{2}}$, $d_4 = \lambda \cdot (c-l) \cdot \sqrt{\frac{\log(2/\delta)}{2}}$, $d_5 = 2 \cdot \lambda \cdot (c-l)$, with the convention that l correspond to 1, c corresponds to 2, r corresponds to 3, (l, c) corresponds to 4 and (c, r) corresponds to 5, the correspondence between times is given by $s = t - t_{\tau_t-1}$, and, for each $s \in \{0, 1, 2, \dots\}$, $m_1(s) = n_{s+t_{\tau_t-1}}(l)$, $m_2(s) = n_{s+t_{\tau_t-1}}(c)$, $m_3(s) = n_{s+t_{\tau_t-1}}(r)$, $m_4(s) = \sum_{i \leq s+t_{\tau_t-1}} \mathbf{1}(x_i \in (l, c))$, $m_5(s) = \sum_{i \leq s+t_{\tau_t-1}} \mathbf{1}(x_i \in (c, r))$. With these conventions, in Lemma 5 we have that $D \leq 9 \cdot \sqrt{5} \cdot (\sqrt{2 \log(2/\delta)} + 4)$ and, for example (the other cases can be proved analogously)

$$\begin{aligned} |J_t(l_{\tau_t}, r_{\tau_t})| &\leq 2 \cdot (\Gamma_t(r) + \Gamma_t(l) + \Gamma_t(l, c) + \Gamma_t(c, r)) \leq 2 \cdot 4 \cdot \max_{k \in [5]} f_k(m_k(s)) \\ &\leq 8 \cdot \frac{D}{\sqrt{t - t_{\tau_t-1} - 5}} \leq \frac{\tilde{c}_\delta}{\sqrt{2}} \cdot \frac{1}{\sqrt{t - t_{\tau_t-1} - 5}} \leq \frac{\tilde{c}_\delta}{\sqrt{t - t_{\tau_t-1}}} \end{aligned}$$

where in the last inequality we used the fact that $t - t_{\tau_t-1} \geq 10$. $\square$

B.3.3 Theorem 4 (Completing the proof)

Proof of Theorem 4. Define τ_T as the last epoch, $t_0 = 0$ and (if not already defined) $t_{\tau_T} = T$.

Due to Lemma 1, we may (and do!) assume that for each $t \in \{1, \dots, T\}$ it holds

$$(\mathbf{U}(c_{\tau_t}) - \mathbf{U}(l_{\tau_t}) \in J_t(l_{\tau_t}, c_{\tau_t})) \wedge (\mathbf{U}(r_{\tau_t}) - \mathbf{U}(c_{\tau_t}) \in J_t(c_{\tau_t}, l_{\tau_t})) \wedge (\mathbf{U}(r_{\tau_t}) - \mathbf{U}(l_{\tau_t}) \in J_t(r_{\tau_t}, l_{\tau_t})).$$

This is because, given our choice $\delta = \frac{1}{T^{5/2}}$, assuming these conditions costs us in the expected regret a further additive term which is no greater than $T \cdot \tilde{C} \cdot T^2 \cdot \delta = \tilde{C} \cdot \sqrt{T}$.

Under these assumptions, notice that for each $\tau \in [\tau_T]$ we have that $x^* \in I_\tau$. In fact, if the confidence intervals are guaranteed to contain the corresponding differences in the expected welfare, every time Algorithm 2 shrinks the active interval is because all the discarded points are guaranteed to be suboptimal.

For each epoch $\tau \in \{1, \dots, \tau_T\}$, define

$$B_\tau := (t_\tau - 1) - t_{\tau-1}.$$

Now, for each epoch $\tau \in \{1, \dots, \tau_T\}$ if $B_\tau \geq \tilde{n}$, then

$$\max_{x \in [l_\tau, r_\tau]} (\mathbf{U}(x^*) - \mathbf{U}(x)) \leq 2 \cdot \tilde{c}_\delta \cdot \sqrt{\frac{1}{B_\tau}}.$$

In fact, assume that $x^* > r_\tau$ (the other cases have similar proofs). Then, leveraging concavity, and recalling that $\inf(J_{t_\tau-1}(l_\tau, r_\tau)) < 0$ and that $x^* \in I_\tau$ (which implies $\frac{x^*-l_\tau}{r_\tau-l_\tau} \leq 2$), we have

$$\begin{aligned} \max_{x \in [l_\tau, r_\tau]} (\mathbf{U}(x^*) - \mathbf{U}(x)) &= \mathbf{U}(x^*) - \mathbf{U}(l_\tau) = \frac{\mathbf{U}(x^*) - \mathbf{U}(r_\tau)}{x^* - r_\tau} (x^* - r_\tau) + \mathbf{U}(r_\tau) - \mathbf{U}(l_\tau) \\ &\leq \frac{\mathbf{U}(r_\tau) - \mathbf{U}(l_\tau)}{r_\tau - l_\tau} (x^* - r_\tau) + \mathbf{U}(r_\tau) - \mathbf{U}(l_\tau) = \frac{x^* - l_\tau}{r_\tau - l_\tau} \cdot (\mathbf{U}(r_\tau) - \mathbf{U}(l_\tau)) \\ &\leq 2 \cdot (\mathbf{U}(r_\tau) - \mathbf{U}(l_\tau)) \leq 2 \cdot \sup(J_{t_\tau-1}(l_\tau, r_\tau)) \leq 2 \cdot |J_{t_\tau-1}(l_\tau, r_\tau)| \\ &\leq 2 \cdot \tilde{c}_\delta \cdot \sqrt{\frac{1}{B_\tau}}, \end{aligned}$$

where the final inequality follows by Lemma [2](#)

Let τ^* be the first epoch from which it holds $x^* \in [l_\tau, r_\tau]$. If $\tau^* \geq 2$, then for each $\tau \in \{2, \dots, \tau^* - 1\}$ it holds that

$$\max_{x \in [l_\tau, r_\tau]} (\mathbf{U}(x^*) - \mathbf{U}(x)) \leq \frac{3}{4} \cdot \max_{x \in [l_{\tau-1}, r_{\tau-1}]} (\mathbf{U}(x^*) - \mathbf{U}(x)).$$

In fact, either for all $\tau \in \{1, \dots, \tau^* - 1\}$ it holds that $r_\tau < x^*$, or for all $\tau \in \{1, \dots, \tau^* - 1\}$ it holds that $l_\tau > x^*$. In the first case, for all $\tau \in \{1, \dots, \tau^* - 1\}$, leveraging concavity and recalling that $x^* \in I_\tau$ (which implies $\frac{x^*-l_\tau}{x^*-l_{\tau-1}} \leq \frac{3}{4}$), we have

$$\begin{aligned} \max_{x \in [l_\tau, r_\tau]} (\mathbf{U}(x^*) - \mathbf{U}(x)) &= \mathbf{U}(x^*) - \mathbf{U}(l_\tau) = \frac{\mathbf{U}(x^*) - \mathbf{U}(l_\tau)}{x^* - l_\tau} \cdot (x^* - l_\tau) \\ &\leq \frac{\mathbf{U}(x^*) - \mathbf{U}(l_{\tau-1})}{x^* - l_{\tau-1}} \cdot (x^* - l_\tau) \\ &\leq \frac{3}{4} \cdot (\mathbf{U}(x^*) - \mathbf{U}(l_{\tau-1})) \\ &= \frac{3}{4} \cdot \max_{x \in [l_{\tau-1}, r_{\tau-1}]} (\mathbf{U}(x^*) - \mathbf{U}(x)), \end{aligned}$$

while the second case can be deduced analogously.

For each $m \in \mathbb{N}$, let $A_m := \{x \in (0, 1) : \exists k \in \{1, \dots, 2^m - 1\}, x = k/2^m\}$ be the

dyadic mesh in $(0, 1)$ of index m . For any epoch $\tau \in \mathbb{N}$, let $m_\tau := -\log_2(c_\tau - l_\tau)$ be the index of the dyadic mesh in $(0, 1)$ at epoch τ of Algorithm 2 (note that $m_\tau \geq 2$ for all $\tau \in \mathbb{N}$ because Algorithm 2 begins with a step-size of $1/4$).

Let $m^* := \min\{m \in \mathbb{N} : \#(A_m \cap (0, x^*]) \geq 4 \text{ and } \#(A_m \cap [x^*, 1)) \geq 4\}$ be the smallest index of the dyadic mesh in $(0, 1)$ such that there are at least 4 points of the dyadic mesh in $(0, 1)$ to the right and to the left of x^* . For each $m \geq m^*$ let $x_1^m < x_2^m < x_3^m < x_4^m \leq x^*$ be the four points of $A_m \cap (0, x^*]$ closest to x^* and $x^* \leq x_5^m < x_6^m < x_7^m < x_8^m$ be the four points of $A_m \cap [x^*, 1)$ closest to x^* . Observe that, for all epochs $\tau \geq \tau^* + 3$, Algorithm 2 selects policies only in the closed interval $[x_1^{m_\tau}, x_8^{m_\tau}]$. Observe further that, for each $m \geq m^* + 1$, it holds

$$\max_{x \in [x_1^m, x_8^m]} (U(x^*) - U(x)) \leq \frac{4}{7} \cdot \max_{x \in [x_1^{m-1}, x_8^{m-1}]} (U(x^*) - U(x)) .$$

In fact, either $\max_{x \in [x_1^m, x_8^m]} (U(x^*) - U(x)) = U(x^*) - U(x_1^m)$ or $\max_{x \in [x_1^m, x_8^m]} (U(x^*) - U(x)) = U(x^*) - U(x_8^m)$. In the first case, leveraging concavity and observing that $\frac{x^* - x_1^m}{x^* - x_1^{m-1}} \leq \frac{4}{7}$, we have

$$\begin{aligned} \max_{x \in [x_1^m, x_8^m]} (U(x^*) - U(x)) &= U(x^*) - U(x_1^m) = \frac{U(x^*) - U(x_1^m)}{x^* - x_1^m} \cdot (x^* - x_1^m) \\ &\leq \frac{U(x^*) - U(x_1^{m-1})}{x^* - x_1^{m-1}} \cdot (x^* - x_1^m) \leq \frac{4}{7} \cdot (U(x^*) - U(x_1^{m-1})) \\ &\leq \frac{4}{7} \cdot \max_{x \in [x_1^{m-1}, x_8^{m-1}]} (U(x^*) - U(x)) . \end{aligned}$$

The second case can be worked out similarly.

Define $\tau^\# := \lfloor 4 + 2 \log_{4/3}(\sqrt{T}) \rfloor$ so that

$$\left(\frac{3}{4}\right)^{\lfloor \frac{\tau^\# - 1}{2} \rfloor} = \left(\frac{3}{4}\right)^{\left\lfloor \frac{\lfloor 4 + 2 \log_{4/3}(\sqrt{T}) \rfloor - 1}{2} \right\rfloor} \leq \left(\frac{3}{4}\right)^{\log_{4/3}(\sqrt{T})} = \frac{1}{\sqrt{T}} .$$

Assume that $\tau^\# < \tau^*$ and $\tau^* + 2 + \tau^\# < \tau_T$ (the other cases can be treated analogously, omitting terms which are not there anymore). Then, the expected regret can

be decomposed as follows:

$$\begin{aligned} \sum_{t=1}^T (U(x^*) - U(x_t)) &= \sum_{\tau=1}^{\tau^\#} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)) + \sum_{\tau=\tau^\#+1}^{\tau^*-1} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)) \\ &+ \sum_{\tau=\tau^*}^{\tau^*+2} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)) + \sum_{\tau=\tau^*+3}^{\tau^*+2+\tau^\#} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)) + \sum_{\tau=\tau^*+3+\tau^\#}^{\tau_T} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)). \end{aligned}$$

We analyze these five terms individually.

For the first one, we further split the sum into two terms, depending on whether or not $B_\tau := t_\tau - 1 - t_{\tau-1} \geq \tilde{n}$. Recalling that for each $\tau \in \{1, \dots, \tau_T\}$ and for each $t \in \{t_{\tau-1} + 1, \dots, t_\tau\}$ Algorithm [2](#) selects the policy x_t in the closed interval $[l_\tau, r_\tau]$, we have that

$$\begin{aligned} \sum_{\substack{\tau=1 \\ B_\tau \geq \tilde{n}}}^{\tau^\#} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)) &\leq \sum_{\substack{\tau=1 \\ B_\tau \geq \tilde{n}}}^{\tau^\#} (B_\tau + 1) \cdot \max_{x \in [l_\tau, r_\tau]} (U(x^*) - U(x)) \\ &\leq \sum_{\substack{\tau=1 \\ B_\tau \geq \tilde{n}}}^{\tau^\#} (B_\tau + 1) \cdot 2 \cdot \tilde{c}_\delta \cdot \sqrt{\frac{\log(2/\delta)}{B_\tau}} \\ &\leq 4 \cdot \tilde{c}_\delta \cdot \sum_{\substack{\tau=1 \\ B_\tau \geq \tilde{n}}}^{\tau^\#} \sqrt{B_\tau} \leq 4 \cdot \tilde{c}_\delta \cdot \tau^\# \cdot \sqrt{T}. \end{aligned}$$

On the other hand, we also have that

$$\sum_{\substack{\tau=1 \\ B_\tau \leq (\tilde{n}-1)}}^{\tau^\#} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)) \leq (\tilde{n} - 1) \sum_{\tau=0}^{\infty} (3/4)^\tau = 4 \cdot (\tilde{n} - 1).$$

Thus, the first term is upper bounded by $4 \cdot \tilde{c}_\delta \cdot \tau^\# \cdot \sqrt{T} + 4 \cdot (\tilde{n} - 1)$.

For the second term, leveraging the definition of $\tau^\#$, we obtain

$$\begin{aligned} \sum_{\tau=\tau^\#+1}^{\tau^*-1} \sum_{t=t_{\tau-1}+1}^{t_\tau} (U(x^*) - U(x_t)) &\leq \sum_{\tau=\tau^\#+1}^{\tau^*-1} \sum_{t=t_{\tau-1}+1}^{t_\tau} (3/4)^{\tau-1} \leq (3/4)^{\tau^\#-1} \cdot \sum_{\tau=\tau^\#+1}^{\tau^*-1} \sum_{t=t_{\tau-1}+1}^{t_\tau} 1 \\ &\leq (3/4)^{\lfloor \frac{\tau^\#-1}{2} \rfloor} \cdot \sum_{\tau=\tau^\#+1}^{\tau^*-1} \sum_{t=t_{\tau-1}+1}^{t_\tau} 1 \leq \sqrt{T}. \end{aligned}$$

For the third term, we further split the sum into two terms, depending on whether or not $B_\tau \geq \tilde{n}$. Proceeding exactly as for the first term, we obtain

$$\sum_{\tau=\tau^*}^{\tau^*+2} \sum_{t=t_{\tau-1}+1}^{t_\tau} (\mathbf{U}(x^*) - \mathbf{U}(x_t)) \leq 3 \cdot 4 \cdot \tilde{c}_\delta \cdot \sqrt{T} + 3 \cdot (\tilde{n} - 1).$$

For the fourth term, we split again the sum into two terms, depending on whether or not $B_\tau \geq \tilde{n}$. If $B_\tau \geq \tilde{n}$, proceeding exactly as for the corresponding part of the first term, we obtain

$$\sum_{\substack{\tau=\tau^*+3 \\ B_\tau \geq \tilde{n}}}^{\tau^*+2+\tau^\#} \sum_{t=t_{\tau-1}+1}^{t_\tau} (\mathbf{U}(x^*) - \mathbf{U}(x_t)) \leq 4 \cdot \tilde{c}_\delta \cdot \tau^\# \cdot \sqrt{T}.$$

Instead, if $B_\tau \leq (\tilde{n} - 1)$, we get

$$\begin{aligned} \sum_{\substack{\tau=\tau^*+3 \\ B_\tau \leq (\tilde{n}-1)}}^{\tau^*+2+\tau^\#} \sum_{t=t_{\tau-1}+1}^{t_\tau} (\mathbf{U}(x^*) - \mathbf{U}(x_t)) &\leq (\tilde{n} - 1) \cdot \sum_{\substack{\tau=\tau^*+3 \\ B_\tau \leq (\tilde{n}-1)}}^{\tau^*+2+\tau^\#} \max_{x \in [l_\tau, r_\tau]} (\mathbf{U}(x^*) - \mathbf{U}(x)) \\ &\leq (\tilde{n} - 1) \cdot \sum_{\substack{\tau=\tau^*+3 \\ B_\tau \leq (\tilde{n}-1)}}^{\tau^*+2+\tau^\#} \max_{x \in [x_1^{m_\tau}, x_8^{m_\tau}]} (\mathbf{U}(x^*) - \mathbf{U}(x)) \\ &\leq 2 \cdot (\tilde{n} - 1) \cdot \sum_{\tau=0}^{\infty} (4/7)^\tau \\ &\leq \frac{14}{3} \cdot (\tilde{n} - 1). \end{aligned}$$

For the last term, we have

$$\begin{aligned} \sum_{\tau=\tau^*+3+\tau^\#}^{\tau_T} \sum_{t=t_{\tau-1}+1}^{t_\tau} (\mathbf{U}(x^*) - \mathbf{U}(x_t)) &\leq \sum_{\tau=\tau^*+3+\tau^\#}^{\tau_T} \sum_{t=t_{\tau-1}+1}^{t_\tau} \max_{x \in [x_1^{m_\tau}, x_8^{m_\tau}]} (\mathbf{U}(x^*) - \mathbf{U}(x)) \\ &\leq \sum_{\tau=\tau^*+3+\tau^\#}^{\tau_T} \sum_{t=t_{\tau-1}+1}^{t_\tau} (4/7)^{\lfloor \frac{\tau-(\tau^*+3)-1}{2} \rfloor} \\ &\leq (3/4)^{\lfloor \frac{\tau^\#-1}{2} \rfloor} \sum_{\tau=\tau^*+3+\tau^\#}^{\tau_T} \sum_{t=t_{\tau-1}+1}^{t_\tau} 1 \leq \sqrt{T}. \end{aligned}$$

Putting everything together, and recalling the definition of $\tau^\#$, the conclusion follows.

□

B.4 Theorem 5 (Upper bound on regret of Tempered Exp3 for Optimal Income Taxation)

We prove this result by **reduction to our baseline model**, as analyzed in Section 3. Assume that $\mathcal{W} = \{w^1, \dots, w^H\}$ with $0 = w^1 < w^2 < \dots < w^H \leq 1$. For each tax bracket $[w^h, w^{h+1})$, *Tempered Exp3 for Optimal Income Taxation* essentially reduces to a separate instance of *Tempered Exp3 for Social Welfare*. Denote

$$\begin{aligned} U_i^h(\mathbf{x}(\cdot)) &= U_i(\mathbf{x}(\cdot)) \cdot \mathbf{1}(\lfloor w_i \rfloor = w^h), & \mathbb{U}_i^h(\mathbf{x}(\cdot)) &= \sum_{j \leq i} U_j^h(\mathbf{x}(\cdot)), \\ \mathbb{U}_i^h &= \sum_{j \leq i} U_j^h(\mathbf{x}_j(\cdot)), & T^h &= \sum_{i \leq T} \mathbf{1}(\lfloor w_i \rfloor = w^h), \\ \mathcal{R}_T^h &= \sup_{\mathbf{x}(\cdot) \in \mathcal{X}_{\mathcal{W}}} E \left[\mathbb{U}_T^h(\mathbf{x}(\cdot)) - \mathbb{U}_T^h \left| \{v_i\}_{i=1}^T, \{w_i\}_{i=1}^T \right. \right]. \end{aligned}$$

It is immediate that

$$\mathcal{R}_T = \sum_h \mathcal{R}_T^h, \text{ and } T = \sum_h T^h.$$

Assume for a moment that the **upper bound** on regret of Theorem 2 (with λ replaced by 1) **applies to each instance** (tax bracket) h , separately. That is, assume that

$$\mathcal{R}_T^h \leq \left(\gamma + \eta \cdot (e - 2)^{\frac{K+1}{K}} \cdot \left(\frac{2K+1}{6} + \frac{1}{\gamma} \right) + \frac{1}{K} \right) \cdot \mathbf{T}^h + \frac{\log(K+1)}{\eta}.$$

Then it follows that

$$\mathcal{R}_T \leq \left(\gamma + \eta \cdot (e - 2)^{\frac{K+1}{K}} \cdot \left(\frac{2K+1}{6} + \frac{1}{\gamma} \right) + \frac{1}{K} \right) \cdot \mathbf{T} + \frac{\mathbf{H} \cdot \log(K+1)}{\eta},$$

and the claims of Theorem 5 are immediate.

It remains to show that indeed the upper bound on regret of Theorem 2 applies to each instance (tax bracket) h . For any given pair of sequences $\{v_i\}_{i=1}^T, \{w_i\}_{i=1}^T$, consider the subsequence of observations i for which $\lfloor w_i \rfloor = w^h$. Along this subsequence, the policy choice reduces to the choice of a tax rate $x_i = \mathbf{x}_i(w^h) \in \mathcal{X}$, and the algorithm *Tempered Exp3 for Optimal Income Taxation* reduces to an instance of the algorithm *Tempered Exp3 for Social Welfare*, with the following **modifications**:

1. Estimated demand $\widehat{G}_i(x, w^h)$ is multiplied by an additional factor $w_i \in [0, 1]$.
2. Estimated social welfare $\widehat{U}_{i+1}(x, w^h)$ is updated with a term for private welfare that includes a time-varying welfare weight $\omega(w_i) \leq 1$, rather than a fixed weight λ .

We need to verify that, with these modifications, the following key claims in the proof of Theorem [2](#) continue to hold:

1. Unbiasedness: $\widehat{U}_i(x, w^h)$ is an unbiased estimator of $\widetilde{U}_i(x, w^h)$, for a suitably discretized version of cumulative social welfare. (Step 2 of the original proof.) In the present setting, discretization requires substituting $\tilde{v}_i$ for v_i , where $\tilde{v}_i = \min\{x \in \mathcal{X} : w_i(1 - x) \geq v_i\}$.
2. Bounded support: $\widehat{U}_i(x, w^h) < \frac{K+1}{\gamma}$, where

$$\widehat{U}_i(x, w^h) = x \cdot \widehat{G}_i(x, w^h) + \frac{\omega(w_i)}{K} \cdot \sum_{x' \in \mathcal{X}, x' > x} \widehat{G}_i(x', w^h).$$

(Step 4 of the original proof.)

3. Bounded second moment of $\widehat{U}_i(x, w^h)$:

$$E_i \left[\widehat{U}_i(x, w^h)^2 \right] \leq \frac{x^2}{p_i(x|w^h)} + \left(\frac{1}{K} \right)^2 \cdot \sum_{x' \in \mathcal{X}, x' > x} \frac{1}{p_i(x'|w^h)},$$

(Step 6 of the original proof.)

Unbiasedness follows as before. To show bounded support, as well as the bound on the second moment, note that we can rewrite

$$\widehat{U}_i(x, w^h) = \left(x \cdot \mathbf{1}(x_i = x) + \frac{\omega(w_i)}{K} \cdot \mathbf{1}(x_i > x) \right) \cdot \frac{y_i \cdot w_i}{p_i(x_i|w^h)}.$$

Recall that $x, \omega(w_i), w_i$, and y_i are all bounded above by 1, and that $p_i(x|w^h) \geq \frac{\gamma}{K+1}$. Bounded support and the bound on the second moment follow. The remaining steps of the proof are as before.