
Stateful Offline Contextual Policy Evaluation and Learning

Nathan Kallus

Cornell University/Cornell Tech

Angela Zhou

UC Berkeley

Abstract

We study off-policy evaluation and learning from sequential data in a structured class of Markov decision processes that arise from repeated interactions with an exogenous sequence of arrivals with contexts, which generate unknown individual-level responses to agent actions that induce known transitions. This is a relevant model, for example, for dynamic personalized pricing and other operations management problems in the presence of potentially high-dimensional user types. The individual-level response is not causally affected by the state variable. In this setting, we adapt doubly-robust estimation in the single-timestep setting to the sequential setting so that a state-dependent policy can be learned even from a single timestep’s worth of data. We introduce a *marginal MDP* model and study an algorithm for off-policy learning, which can be viewed as fitted value iteration in the marginal MDP. We also provide structural results on when errors in the response model leads to the persistence, rather than attenuation, of error over time. In simulations, we show that the advantages of doubly-robust estimation in the single time-step setting, via unbiased and lower-variance estimation, can directly translate to improved out-of-sample policy performance. This structure-specific analysis sheds light on the underlying structure on a class of problems, operations research/management problems, often heralded as a real-world domain for offline RL, which are in fact qualitatively easier.

1 Introduction

Offline reinforcement learning seeks to reuse existing data to evaluate and learn novel policies and is crucial in applications with limited freedom to experiment but plentiful logged data. In general Markov decision processes (MDPs), offline reinforcement learning can be very difficult, as we must understand the effect of actions in each state and time, whether in model-based (e.g., learn the transition kernel) or model-free methods (e.g., learn Q -functions). However, many practically-relevant problems fit in simpler, more tractable classes of MDPs with “sequential decision-making” but not “longitudinal data”, for example because transitions arise in a stochastic system from exogenous arrivals. In this paper, we study off-policy evaluation and optimization from observational data in this special class. At each timestep the *same* contextual response model generates both transitions and rewards. The setting is a variant of offline contextual bandits with constraints, where the same randomness generates transitions in the system state (status of the constraints) and rewards in the system. We call this setting, common in operations management, “stateful” to emphasize the well-understood and simple system state, like inventory state, in contrast to the unknown potentially high-dimensional “contextual” response model, like an individual’s propensity to purchase, that must be learned.

We first describe some stylized examples to illustrate how previously studied classical problems in fact share this broader structure: consider personalized dynamic pricing with inventory constraints, or managing a rideshare system and repositioning vehicles by making price offers to individuals. The system state includes capacities of each resource, or locations of cars in the system. Individuals with contexts (covariates) arrive exogenously. The system takes actions, such as personalized price or trip offers. Given a context and action, the individual response changes system dynamics: the purchase of a product consumes resources, or accepting a price offer and ride from one location to another moves cars. But given that we can offer the resource at all, the state of the system does not further affect

the response except by affecting our pricing decisions.

We focus on the evaluation and optimization of state-dependent policies from offline trajectories collected from these system dynamics. Confounding is of particular concern in such observational data. Naturally, system actions can be spuriously correlated with outcomes. For example, we expect observational operational data to bias toward higher-revenue actions: higher price offers are made to individuals deemed more likely to accept. However, in our setting, the underlying system state does not causally affect the individual response and therefore, unlike individual-level covariates, the system state is not a confounder, making it easier to learn the response model from observational data, and then use it to design sequential policies. Ultimately we show how this special structure of problems, common to operations problems arising from uniformized stochastic systems permits developing specialized OPE.

The contributions of this paper are as follows. We model the above structure and study estimation of the transition probabilities in a lifted *marginal* MDP via single timestep off-policy evaluation. Therefore we reduce the analysis of a high-dimensional, continuous state space MDP to the standard tabular analysis of our marginal MDP, reducing the number of nuisances from $T + 1$ for doubly-robust OPE to *just two* in our setting. We show $\tilde{O}(T^3/\epsilon^2)$ trajectories are required for off-policy optimization to achieve ϵ -suboptimal value, where T is the horizon (omitting logarithmic terms). We study error amplification for the dynamic and capacitated pricing example and show that bias from naive model-based approaches would generally persist in realistic scenarios. We validate the theory and structural analysis in simulations where we improve on naive model-based approaches and generic offline RL.

2 Problem setup: Stateful off-policy evaluation and learning

We first describe the generic full-information MDP that generates our data before describing the restrictions that characterize the stateful setting. For ease of reference, we partition the state space of the full-information MDP into a product space of the *discrete system state space* $\mathcal{S}$, potentially continuous *context/covariate space* $\mathcal{X}$, and discrete covariate-conditional response space $\mathcal{Y}$: $\mathcal{S} \times \mathcal{X} \times \mathcal{Y}$. The inclusion of Y in the state variable is purely informational. Consider a finite-horizon setting with $T + 1$ timesteps, and denote the initial system state s_0 ; timesteps are indexed $0, \dots, T$. Uppercase (S, X) indicates random variables; lowercase (s, x) fixed values; and prime (s', x') next-timestep values. Let $\mathcal{A}(s, x)$ denote the

discrete action space feasible from the state (s, x) . A contextual policy $\pi_t: \mathcal{S} \times \mathcal{X} \mapsto \Delta^{\mathcal{A}}$ maps from system state/context to a distribution over actions, where $\Delta^{\mathcal{A}}$ is the set of distributions defined on $\mathcal{A}$, so that $\pi_t(a | s, x) = \mathbb{P}(A = a | S = s, X = x)$ gives the probability of taking action a given state and context information. Let $\pi = \{\pi_t\}_{t=0, \dots, T}$ denote the MDP policy in a function class $\Pi_{0:T}$. Reward is a known deterministic function of next state transition, $R(s, a, s')$.

“Stateful contextual” structure. We next specify the restrictions on this MDP that give rise to our “stateful” setting. These are illustrated in Figure 1. Roughly: contexts arrive exogenously and contextual responses Y come from a stationary conditional distribution $\mathbb{P}(Y | X, A)$ and deterministically generate the system transitions. We henceforth use the shorthand $P(s', x', y | s, x, y_{-1}, a)$ for the transition model under this convention, although dependence on y_{-1} is purely artificial/notational and will be omitted in general.

We formalize as assumptions the general structure that appears commonly in more specialized problem contexts elsewhere that describes the “stateful” setting.

Assumption 1 (Exogenous context process). The transition factorizes as

$$P(s', x', y | s, x, y_{-1}, a) = P(s', y | s, x, a) f(x') \\ \forall s, s', x, x', y_{-1}, y, a$$

Assumption 2 (Contextual-response transitions). We know $s'(s, y): \mathcal{S} \times \mathcal{Y} \mapsto \mathcal{S}$ such that when s' is not absorbing from s we have:

$$P(s', y | s, x, a) = \delta_{s'-s'(s, y)} P(Y = y | x, a)$$

We can easily extend to random transitions given responses, but focus on deterministic for concreteness and as it captures the most relevant application settings. Assumption 1 arises from contextual bandits or uniformizing (with contexts) a stochastic system [Gallego et al. (2019); Meyn and Meyn (2008)]. Assumption 2 reflects the offline contextual bandit nature of the problem and encodes that Y_t is independent of the originating state. [El Shar and Jiang (2020)] leverages a factorization with exogenous information, but not a contextual response model and notes that the “system transition function” construction is the norm in control/operations research [Bertsekas and Tsitsiklis (1996); Powell (2007)].

For ease of presentation we introduce $\{s', y | s, a\}$ as the pairs of next states and contextual responses reachable from s .

Definition 1 (Reachable state transition-potential

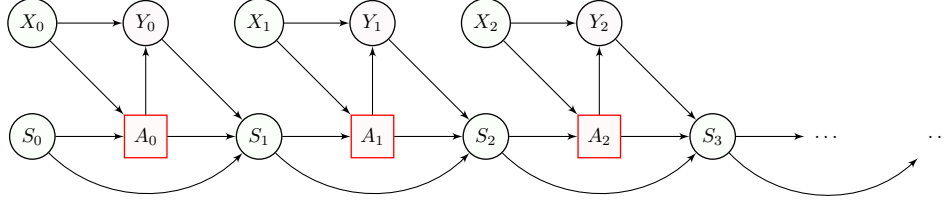


Figure 1: The “stateful” decision model we consider; rewards are functions of S_t, A_t, S_{t+1} .

outcomes).

$$\begin{aligned} \{(s', y) \mid s\} &:= \\ \{(s', y) \in \mathcal{S} \times \mathcal{Y} : \exists x, a \text{ s.t. } P(s', y \mid s, x, a) > 0\} \end{aligned}$$

The corresponding full-information MDP is $\mathcal{M} = (\mathcal{S} \times \mathcal{X} \times \mathcal{Y}, \mathcal{A}, P, R, T)$ where P is the full-information transition kernel. Our *observational* data comprises of N trajectories; denote individual observations as S_t^i for t timestep of trajectory i : $\{S_t^i, X_t^i, A_t^i, Y_t^i, R_t^i\}_{0:T}^n$.

Without loss of generality we omit the Y_{t-1} information state from the full-information *value/reward to go* V - and state-action value Q -function, since $V_t(s, x, y_{-1}) = V_t(s, x)$:

$$\begin{aligned} V_t^\pi(s, x) &= \mathbb{E}_\pi \left[\sum_{t'=t}^T R_{t'} \mid S_t = s, X_t = x \right], \\ Q_t^\pi(s, x, a) &= \mathbb{E} \left[\sum_{t'=t}^T R_{t'} \mid S_t = s, X_t = x, A_t = a \right] \end{aligned}$$

It is useful to define the *context-marginalized* value function $\tilde{V}_t^\pi(s) = \mathbb{E}[V_t^\pi(s, X)]$, the value function at system state S_t marginalized over context distribution X_t , and analogously $\tilde{Q}$. Under assumptions 1 and 2 and the notation of defn. 1, the Q -function in the full-information MDP is:

$$Q_t^\pi(s, x, a) = \sum_{(s', y) \mid s} \mathbb{P}(Y = y \mid x, a) (R(s, a, s') + \tilde{V}_{t+1}^\pi(s')) \quad (1)$$

Finally, throughout the paper we will assume the behavior policy is *stationary*, and *not* history-adapted for a simpler statement of our results.¹

Assumption 3 (Stationary behavior policy). $\pi_t^b(a \mid s, x)$ is stationary: possibly time-varying, but not history-dependent upon $\{S_{t'}, X_{t'}, A_{t'}, Y_{t'}\}_{t' < t}$.

Specific examples of stateful problems. We discuss illustrative examples. The first example, single-item personalized and dynamic pricing with inventory

constraints, is adapted from classical models for network revenue management (Gallego et al. (2019); Gallego and Van Ryzin (1997)).²

Example 1 (Single item personalized and dynamic pricing). $Y_t \in \{0, 1\}$ is purchase/no-purchase, respectively, and $A_t \in \{0, 1\}$ is whether a discount of value d is not or is offered. Let $p(a)$ be the price corresponding to taking action $A = a$. The reward is fixed given transition to s' : $R(s, a, s') = p(a)\mathbb{I}[y = 1]\mathbb{I}[s > 0]$. For short let $R(a)$ denote price of product under a , e.g. reward only received if item is sold, and we can only sell if we have stock so $s'(s, y) = \mathbb{I}[s > 0, y = 1](s - 1) + \mathbb{I}[s > 0, y = 0]s$. Denote the difference of value functions as $\Delta V_t^\pi(s) = \tilde{V}_t^\pi(s - 1) - \tilde{V}_t^\pi(s)$, then the full-information Q_t function is 0 for $t = T$, $\mathbb{P}(Y = y \mid x, a)R(a)\mathbb{I}[s > 0]$ for $t = T - 1$, and for $t < T - 1$:

$$\begin{aligned} Q_t(s, x, a) &= \\ \mathbb{P}(Y = y \mid x, a) (R(a)\mathbb{I}[s > 0] + \Delta V_{t+1}^\pi(s)) + \tilde{V}_{t+1}^\pi(s). \end{aligned}$$

Next we describe in words other examples that also fit in the model of Figure 1 but defer their specific mathematical formulations to the appendix.

Example 2 (Multi-item network revenue management (informal) Gallego and Van Ryzin (1997); Gallego et al. (2019)). This extends Example 1 with multivariate outcomes (contextual demands for different products). We augment the exogenous context arrival process with *product arrival types* and product-conditional context distributions.

Example 3 (Spatial pricing and repositioning (informal, contextual adaptation of El Shar and Jiang (2020); Bimpikis et al. (2019))). The state space is the number of cars at each station in a ridesharing system. We augment the exogenous information process via uniformized arrivals at a station and origin-destination requests. The individual contextual response is ride acceptance/rejection at a price; reward is revenue and a lost sales cost.

¹We leave finer-grained analysis of dependent data, where the benefits of pooling data also trade-off against dependence and mixing rates, for future work; Corollary 2 highlights how this assumption may be removed via standard arguments.

²Instead of assuming arrivals would deterministically purchase and setting the decision variable to be fare availability, we consider a stochastic demand response model.

Structure that satisfies or does not satisfy assumptions. We have discussed classical examples that instantiate these assumptions. However, more complex modeling could violate them. Assumption 1 would not be true if customers had full observation of the system and responded to it, such as in queuing if customers can observe queue length and balk. Or, for example 3, if customer arrivals are correlated with system state due to unobserved confounders, such as weather patterns that lead to higher propensity to accept a ride and higher customer demand at other locations. In the context of example 3, Assumption 2 is true if the underlying system state S_t (cars at other stations) is not a confounder because it does not affect an *individual's* demand in response to price. However, if system transitions modeled stochastic travel times where system state variables (such as congestion) did causally affect outcomes, Assumption 2 may not hold.

3 Related work

In the main text we only highlight the most closely related work; see Section C.1 for further discussion.

Off-policy policy learning for offline sequential decision-making. There has been extensive work on off-policy evaluation and learning in the sequential setting. We focus on work that builds on statistical model-free approaches, including doubly robust off-policy evaluation in incorporating value-function control variates [Thomas and Brunskill (2016); Jiang and Li (2016); Zhang et al. (2013)], and study of the efficient influence function [Kallus and Uehara (2019a); Bibaut et al. (2019); Kallus and Uehara (2019b)], as well as MIS or fitted-Q-evaluation [Yin et al. (2021); Duan et al. (2020); Le et al. (2019); Hu et al. (2021)].

In general, off-policy evaluation in the sequential setting either includes rejection sampling on entire trajectories (even with doubly-robust augmentation) [Thomas and Brunskill (2016)], or introduces marginalized density ratios [Yin et al. (2021); Kallus and Uehara (2019a)] which in the finite-horizon setting cannot be optimized in a backwards-recursive fashion or are policy dependent. The latter prevents direct translation of improvements in statistical OPE to off-policy policy optimization except by exhaustive search over the policy class. [Nie et al. (2020)] similarly specializes OPE to a different setting, optimal stopping, which admits policy-independent nuisance functions. We similarly develop structure-dependent improvements in dependence on nuisance functions, but for different structure.

Our estimator, derived via the modeling analysis in the next section, does not require rejection sampling on en-

tire trajectories. Therefore we show statefulness is in fact more closely related to single-timestep off-policy evaluation and learning [Dudik et al. (2014); Kitagawa and Tetenov (2015); Swaminathan and Joachims (2015); Wager and Athey (2017)]. We do not claim novelty relative to the extensively-studied doubly-robust estimation in sequential OPE [Jiang and Li (2016); Tang et al. (2019)]; rather we show that specializing to policy structure allows for retaining statistical improvements from double robustness with reduced dependence on nuisance functions (two instead of $T + 1$).

Online contextual decision-making with constraints and algorithmic analysis under known distributions. Online contextual decision-making with constraints. There is an extensive literature on either contextual or stateful problems in operations research, including online learning. Typically contexts are discrete, known types. We highlight work that studies online learning in constrained systems, such as (episodic) inventory/revenue management, [Huh et al. (2011); Besbes and Zeevi (2012); Agrawal and Jia (2019)], or contextual decisions such as covariate-based dynamic pricing [Cohen et al. (2016); Javanmard and Nazerzadeh (2016); Qiang and Bayati (2016); Shah et al. (2019); Ban and Keskin (2020); Chen et al. (2021)]. These approaches are typically model-based: they require uncontextual demand distributions (known, or learned online) or impose parametric restrictions. Contextual bandits with knapsack (CBwK) does consider both contexts and statefulness. [Badanidiyuru et al. (2018)]³ The closest work is [Agrawal et al. (2016)], which uses single-timestep offline policy optimization but considers the Lagrangian relaxation of the resource constraints: regret guarantees are on the Lagrangian and the policy satisfies constraints in expectation rather than with probability 1.

In contrast to the online setting where completely randomized exploration is possible, we are interested in characterizing the setting of learning a dynamic policy from *offline* off-policy data, *without* the ability to set an exploration policy to collect more information. Relative to CBwK and pricing bandits, we consider a general MDP embedding and our sample complexity analysis and algorithm do not require specific structure of the reward beyond assumptions 1 and 2.

Our approach is particularly beneficial in handling high-dimensional context variables X_t . Naively analyzing approximate linear program arising from state

³We discuss CBwK for a full discussion of related work. But while our framework can readily handle unknown or multiple behavior policies, we do not consider data directly collected from a bandit algorithm (i.e. outcome-adapted data subject to adaptive sequential learning bias).

aggregation on X_t incurs statistical bias in general due to discretization. On the other hand, structure-specific analysis of any such problem, such as network revenue management or online packing problems, will generally obtain stronger approximation guarantees for the online setting although we do not comment on direct translation of online algorithms or approximation algorithm-type guarantees, e.g. benchmarked to the fluid relaxations, to the contextual setting. Note [Bray \(2019\)](#) highlights dependence of recent constant-regret approximation guarantees on discrete distributions, while our setting of contextual responses corresponds to the case of continuous valuations. We defer a comprehensive comparison to future work. Our work is on offline evaluation in such problem settings, and can for example allow off-policy evaluation of policies (when they are not history-adapted) from confounded observational data.

4 Off-policy evaluation and learning in the *Marginal MDP*

Marginal MDP construction. Section [2](#) described the generating process of the data. We now marginalize over contexts and (policy-induced) outcomes in a lifted *marginal MDP* on a discrete state space and continuous action space where actions are given by policy parameters. This MDP is purely a *conceptual device* which is used in the analysis. Direct OPE methods cannot be used in the marginal MDP because observations in the dataset correspond to variation over different *actions*, but not necessarily different *policies* that are actions in the marginal MDP. We develop this construction to justify the use of single-timestep off-policy evaluation, which we denote as $P(Y = y \mid \pi)$:

$$P(Y = y \mid \pi) = \sum_{a \in \mathcal{A}} \mathbb{E}[\pi_e(a \mid X) \mathbb{P}(Y = y \mid a, X)] \quad (2)$$

To summarize, the marginal MDP is $\tilde{\mathcal{M}} = (\mathcal{S}, \Pi, \tilde{P}, \tilde{R}, T)$, where the action space is the space of (s -dependent) policy functions of x and transitions and rewards marginalize over context arrivals. The key modeling insight is that expectations over individual exogenous arrivals may be estimated via a distribution of iid arrivals; e.g. estimate eq. [\(2\)](#) by single-timestep off-policy evaluation.

The marginal MDP state space is the system state space, $\mathcal{S}$.

The action space is the set of parametrized policies, $\mathcal{A}(s) = \Pi(s)$, where $\Pi(s) = \{\pi(s, \cdot) \in \Pi\}$ is the set of policy functions given s .

Transitions between s and $s'(s, y)$ occur with probability $P(Y = y \mid \pi)$, (eq. [\(2\)](#))

Reward is $\tilde{R}(s, \pi) = \sum_a \sum_{s', y \mid s} \mathbb{E}[\pi(a \mid X) \mathbb{I}[Y(a) = y] R]$, the expected reward induced by context-conditional policy actions and corresponding outcomes, where $R = R(s, a, s'(s, y))$.

By construction, policy values and optimal policies are equivalent in the marginal and full-information MDPs (under a policy class that is a product class in s). (Note that higher-order moments are not equivalent.)

Proposition 1. Assume the policy class Π is a product space over $s \in \mathcal{S}, t \in [T]$. The marginal MDP $\tilde{\mathcal{M}} = (\mathcal{S}, \Pi, \tilde{P}, \tilde{R}, T)$ has the same optimal policy, and policy value $V(s), Q(s)$, as the full-information MDP with policy class Π and marginal policy values $\tilde{V}(s), \tilde{Q}(s)$ when $\mathcal{M} = (\mathcal{S} \times \mathcal{X} \times \mathcal{Y}, \mathcal{A}, P, R, T)$.

Example: Marginal value function for single-item pricing. In the marginal MDP for Example [1](#),

$$\begin{aligned} \tilde{P}(s - 1 \mid s, \pi) &= P(Y = 1 \mid \pi) \\ s'(s, y) &= \mathbb{I}[s > 0] \mathbb{I}[y = 1](s - 1) \\ \tilde{V}_t^{\pi_{t:T}}(s) &= \\ \tilde{R}(s, \pi) + \tilde{V}_{t+1}^{\pi_{t+1:T}}(s) + P(Y = y \mid \pi) \Delta V_{t+1}^{\pi_{t+1:T}}(s) \end{aligned}$$

Estimation via fitted value evaluation and iteration in the marginal MDP. We define the propensity score and outcome model as follows:

$$\begin{aligned} e(a \mid x) &= P(A_t = a \mid X = x) \\ \mu(y \mid a, x) &= P(Y = y \mid A = a, X = x). \end{aligned}$$

The propensity score only controls for X : while we allow the underlying behavior policy to be state-dependent, Assumption [2](#) implies that adjusting for X is sufficient to estimate the marginalized transition, eq. [\(2\)](#), because the state does not affect the outcome. To achieve the orthogonality and rate double-robustness benefits of the doubly-robust estimator we next introduce, we use two-fold sample splitting in trajectories and timesteps. We use cross-time fitting and introduce folds that partition trajectories and timesteps $k(i, t)$. For $K = 2$ we consider timesteps interleaved by parity (e.g. odd/even timesteps in the same fold). We let $k(i, t)$ denote that nuisance $\hat{\mu}^{-k(i, t)}$ is learned from $\{X_{t'}^{(i)}, Y_{t'}^{(i)}\}_{i \in \mathcal{I}_{k(i)}, t' \bmod 2 = t \bmod 2}$, e.g. from the $-k(i)$ trajectories and from timesteps of the same evenness or oddness but is only used for evaluation in the other fold. Interleaving between timesteps insures downstream policy evaluation errors are independent of errors in nuisance evaluation at time t .

We let $\hat{P}(Y = y \mid \pi)$ denote the empirical estimate: we verify that the standard doubly robust estimator for single-timestep offline policy learning, reweighting the empirical transitions in observational data, estimates the transition probabilities in the marginal MDP.

Proposition 2 (Single-time-step doubly robust estimator of transitions in the marginal MDP.). Let

$$\Gamma_t^i(y | a) := \frac{\mathbb{I}[Y_t^i=y] - \hat{\mu}^{-k}(y|A_t^i X_t^i)}{\hat{e}_t^{-k}(A_t^i|X_t^i)} \mathbb{I}[A_t^i = a] + \hat{\mu}^{-k}(y | a, X_t^i)$$

$$\hat{P}(Y = y | \pi) := (NT)^{-1} \sum_{t=1}^T \sum_{i=1}^N \sum_{a \in \mathcal{A}} \pi(a | X_t^i) \Gamma_t^i(y | a). \quad (3)$$

$\hat{P}(Y = y | \pi)$ is an unbiased estimator of $P(Y = y | \pi)$ if at least one of $\hat{\mu}$ or $\hat{e}$ are unbiased.

Proposition 2 verifies *orthogonality*, that the estimator is doubly-robust against misspecification of one of μ or e . The estimator only adjusts for contexts because Assumption 2 specifies that the state variable is not a confounder. Proposition 2 considers the stationary case; when X_t is time-varying but non-adversarial with fixed distributions, similar data-pooling is possible by estimating density ratios.⁴

Given a generic estimator $\hat{P}(Y = y | \pi)$ for the marginal transition probability $P(Y = y | \pi)$, we can construct a Q -estimate as follows: $\hat{P}(Y = y | \pi)$ can be the doubly robust estimator as in Proposition 2 or alternatively the IPW estimator (simply let $\hat{\mu}^{k(i,t)} = 0$ in Proposition 2) or direct method estimator (simply let $\hat{e}^{k(i,t)} = \infty$ in Proposition 2). We use backwards recursion to evaluate $\hat{V}_t^{\pi_{t+1:T}}(s)$ using model-based evaluation with $\hat{P}(Y = y | \pi)$ in the marginal MDP.

$$\hat{Q}_t^{\pi, \pi_{t+1:T}}(s, \pi) = \sum_{(s', y) | s} \hat{P}(Y = y | \pi) \left(R(s, a, s') + \hat{V}_{t+1}^{\pi_{t+1:T}}(s') \right). \quad (4)$$

Policy Learning. When the policy space is in fact a product set over the state space (i.e., the policy being optimized can vary independently for every value of the state), we study a policy learning proposal in Algorithm 1 which implements backwards-recursive policy learning (which can be understood as fitted value iteration in the marginal MDP) to determine the optimal policy vector π .

5 Analysis

Sample complexity. We first provide a generalization bound for Algorithm 1 on the out-of-sample regret

⁴In revenue-management settings, it is common for X_t arrivals to be nonstationary. While online algorithms considers adversarial arrival distributions, relevant arrivals may also have highly structured nonstationarity, e.g., “business-class” arrivals arriving later on. To a limited extent, *adversarial* arrivals could also be modeled by robustness, e.g., using the approach of Kallus and Zhou (2020) over density ratios for each timestep’s subproblem in Algorithm 1.

Algorithm 1 Backwards-Recursive Policy Learning

- 1: Input: estimate $\hat{P}(Y = y | \pi)$, policy class $\Pi_{0:T}$
 - 2: **for** $t = T, \dots, 0$ **do**:
 - 3: **for** $s \in \mathcal{S}$ **do**:
 - 4: Estimate off-policy value $\hat{Q}_t^{\pi, \hat{\pi}_{t+1:T}}(s, \pi)$ via eq. (4)
 - 5: Optimize $\hat{\pi}_{t,s}^* \in \arg \max_{\pi \in \Pi_t(s)} \hat{Q}_t^{\pi, \hat{\pi}_{t+1:T}}(s, \pi)$ and update $\hat{V}_t^{\hat{\pi}_{t:T}}(s) \leftarrow \hat{Q}_t^{\hat{\pi}_{t:T}}(s, \hat{\pi}_{t,s}^*)$
 - 6: **return** $\hat{\pi}^* = \{\hat{\pi}_{t,s}^* : t \in [T], s \in \mathcal{S}\}$
-

$\tilde{V}_0^{\hat{\pi}^*}$, the true value achieved by the sample-optimal policy $\hat{\pi}^*$, relative to the best-in-class policy, $\tilde{V}_0^{\pi^*}$. We assume the policy class at a given s, t has restricted functional complexity in the sense of a finite entropy integral of the covering numbers Van Der Vaart and Wellner (1996); Wainwright (2019). In the main text we use the VC dimension d_{vc} for binary actions; in the appendix we include corresponding statements for multi-class notions such as Natarajan dimension Mohri et al. (2018).

Theorem 1 (Sample complexity and rate double-robustness). Suppose $\nu^{-1} \leq e(a | x) \leq 1 - \nu^{-1}$ uniformly over a, x , for $\nu > 0$, (overlap) and for some rates $0 < r_1, r_2 < 1$ and constants C_1, C_2 , we have uniformly consistent estimation of nuisance estimates

$$\mathbb{E}[(\mu(y | a, X) - \hat{\mu}(y | a, X))^2] = o_p(n^{-r_1}),$$

$$\mathbb{E}[(e(a | X) - \hat{e}(a | X))^2] = o_p(n^{-r_2})$$

where $r_1 + r_2 \geq 1$. Then there exists a random variable $\kappa = o_p((NT)^{-\frac{1}{2}})$ so that w.p. $\geq 1 - \delta$,

$$\tilde{V}_0^{\pi^*} - \tilde{V}_0^{\hat{\pi}^*} \leq \frac{5\nu R^{\max}(T^{1/2} + 1/2 T^{3/2}) \sqrt{d_{\text{vc}} \log(\frac{5T|Y|}{\delta})}}{\sqrt{N}} + \kappa.$$

The κ term arises because we decompose the value difference with an *oracle* estimator using the true nuisance functions μ, e , and obtain a high-probability bound on the leading order term. The final bound is of order $O_p(N^{-\frac{1}{2}} T^{\frac{3}{2}})$. The proof follows standard techniques, combining single time-step uniform convergence with the performance difference lemma. However, it is the previous modeling analysis and our derived estimator that permits this reduction.

The main improvement in Theorem 1 is in specializing to statefulness so only two nuisance functions are required, rather than $T \times |\Pi_{0:T}|$ many as would arise in the case of Q -function nuisances. In appendix section C we also discuss improvements in dependence on concentratability coefficients/sequential overlap.

Finally, we verify the nuisance rates are as achievable from pooled episodes as they would be from iid data.

In the main text, $\lesssim$ omits polylogarithmic factors; see appendix section B.3 for a full statement. This is a direct corollary of Bojun (2020) and a convergence rate from mixing data given in (Farahmand and Szepesvári, 2012, Thm. 5); α describes a mild capacity assumption on the covering numbers of nonparametric nuisance function classes.

Corollary 2 (Estimation of nuisance functions). *Suppose nuisance function estimator $\hat{e}(a | x)$ is learned from pooled-episode data $\{(X_t^i, A_t^i, Y_t^i)\}_{0:T}^n$, and $\mu(y | a, x)$ from $\{(X_t^i, A_t^i)\}_{0:T}^n$ by regularized least squares. Assume that A, Y are bounded a.s., and well-specification of e, μ . Assume there exist $C > 0, 0 \leq \alpha < 1$ such that for any $\epsilon, R > 0$, the covering numbers are bounded, $\mathcal{N}_2(\epsilon, \mathcal{B}_R, X_{1:N}) \leq C(R/\epsilon)^{2\alpha}$. Then there exists $c_1, c_2 > 0$ such that for large enough n , for any fixed $\delta \in (0, 1)$, with probability $\geq 1 - \delta$,*

$$\begin{aligned}\mathbb{E}[(e(a | X) - \hat{e}(a | X))^2] &\lesssim n^{-\frac{1}{1+\alpha}}, \\ \mathbb{E}[(\mu(a | X) - \hat{\mu}(a | X))^2] &\lesssim n^{-\frac{1}{1+\alpha}}.\end{aligned}$$

Corollary 2 verifies minimax rate-optimality of estimation in this setting (Farahmand and Szepesvári (2012); Yang and Barron (1999), e.g. estimating the nuisance functions from pooled episodes achieves the minimax nonparametric rate that prevails in the absence of dependence up to polylogarithmic terms.

Remark 1 (Extension to continuous states). We can extend to doubly-robust policy evaluation with continuous transitions by invoking recent advances in estimating counterfactual distributions (Kennedy et al. (2021)) but further work is required for function approximation for policy learning. See Section C.4.

Structural analysis: error propagation in dynamic pricing. We study the possible drawbacks of naively plugging-in a confounded model (DM) by analyzing the sequential error propagation. We provide structural conditions for error *persistence* in the sequential setting: error could “persist” if model error in transitions and value functions continues to impact downstream learned policies or “attenuate” if these cancel out in the sequential setting. We specialize to Example 1; similar results should hold for other settings with less interpretable sufficient conditions. Define the threshold θ^* , true conditional expectation ratio $\Delta\mu^*$, and optimal (threshold) policy π^* :

$$\begin{aligned}\theta_t^*(s, \Delta V) &= \frac{R(0) + \Delta V}{R(1) + \Delta V}, \quad \Delta\mu^*(x) = \frac{\mu(1|x)}{\mu(0|x)}, \quad (5) \\ \pi^*(1 | s, x) &= \mathbb{I}[\Delta\mu^*(x) > \theta_t^*]. \quad (6)\end{aligned}$$

The confounded-optimal policy incurs error from $\Delta\hat{\mu}$ or $\Delta V^{\hat{\pi}^*}$. We reparametrize the decision boundary on $\Delta\hat{\mu}$ relative to $\theta^*, \Delta\mu^*$. Then the biased threshold $\hat{\theta}^*$

is related to the true θ^* by the pointwise error $\delta(a, x)$:

$$\delta(a, x) = \hat{\mu}(y | a, x) - \mu(y | a, x), \quad (7)$$

$$\hat{\theta}^* = \theta_t^*(s) \cdot (1 + \delta(0, x)/\mu(1|0, x)) - \delta(1, x). \quad (8)$$

When self-evident we omit dependence of θ^* on arguments that remain fixed for brevity. Error “persists” if the error at different timesteps, including from value estimation, persists in the same direction relative to the optimal policy. We provide a sufficient condition to conclude the direction of error induced from downstream errors in value estimation. In the main text we state a special case; the appendix includes the full theorem for $t < T - 2$ with less interpretable conditions.

Theorem 3 (Conditions for error persistence). *For $t = T - 2$, assume $R(1) > R(0)$, without loss of generality. Then, for $s > 2$,*

$$\begin{aligned}\mathbb{E}[-\tau(X)\mathbb{I}[\hat{\theta}_{T-1}^* \leq \Delta\mu^*(X) \leq \theta_{T-1}^*]] &\geq 0 \\ \implies \hat{\theta}_{T-2}^*(\Delta V_{T-1}^{\hat{\pi}_{T-1}^*}) &< \hat{\theta}_{T-2}^*(\Delta V_{T-1}^{\pi_{T-1}^*}).\end{aligned}$$

We discuss implications of Theorem 3 for bias persistence in the context of Example 1.

Example 4 (Error persistence in Example 1). Suppose $\delta(1, x) > 0 > \delta(0, x)$ uniformly over x , for example if historical price increases were targeted towards those more likely to purchase them and discounts were targeted to those less likely to purchase overall. Then for any s , $\Delta V, \hat{\theta}^*(s, \Delta V) \leq \theta^*(s, \Delta V)$. By assumption on, e.g. price elasticity so that $\tau(x) \leq 0$, we expect the sufficient condition of Theorem 3 to be true so that bias persists; for $s > 2$,

$$\hat{\theta}_{T-2}^*(\Delta V_{T-1}^{\hat{\pi}_{T-1}^*}) < \hat{\theta}_{T-2}^*(\Delta V_{T-1}^{\pi_{T-1}^*}) < \theta_{T-2}^*(\Delta V_{T-1}^{\pi_{T-1}^*}).$$

6 Empirics

Data-generating process. We consider a simple example based on single-product dynamic pricing (example 1), with a response model that is a Δ -weighted mixture model of a logistic specification and a nonlinear specification, where $\sigma(\beta^\top x) = (1 + \exp(-\beta^\top x))^{-1}$.

$$\begin{aligned}\mu(1 | a, x) &= (1 - \Delta)\sigma(\beta^\top x + \beta_0 p_a) + \Delta\sigma(x_0^2 p_a), \quad (9) \\ e_t(1 | x) &= \sigma(-0.8\beta^\top x) \quad (10)\end{aligned}$$

We generate the data corresponding to the outcome specification for parameters $\beta = [-0.75, 0.75]$, $\beta_0 = -2$. We learn outcome models $\hat{\mu}$ by either logistic regression (for DM, direct method or DR, doubly robust) or a neural network for a nonparametric nuisance estimate (DR-nonpara), and the behavior policy by (well-specified) logistic regression. We

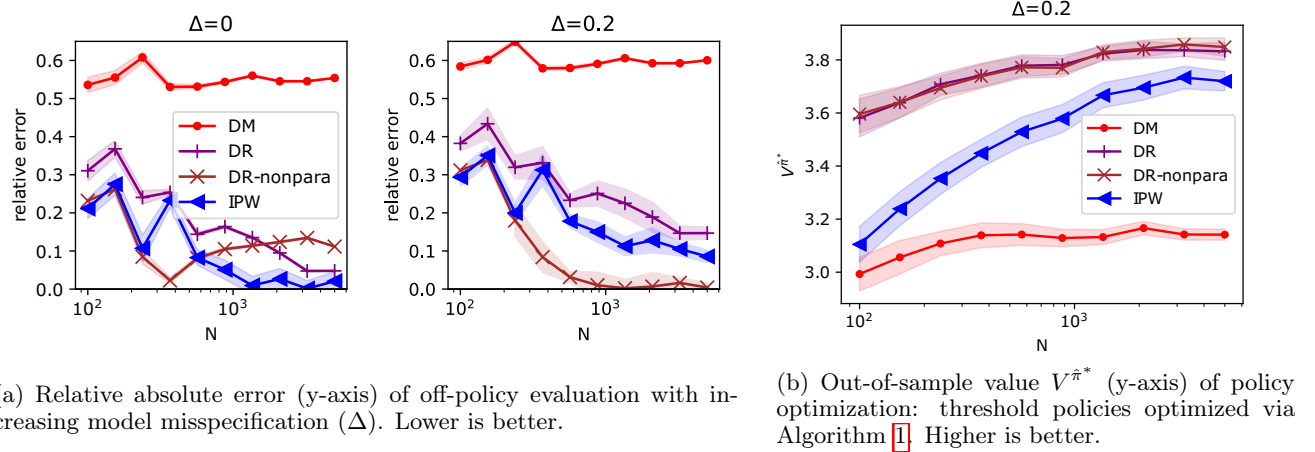
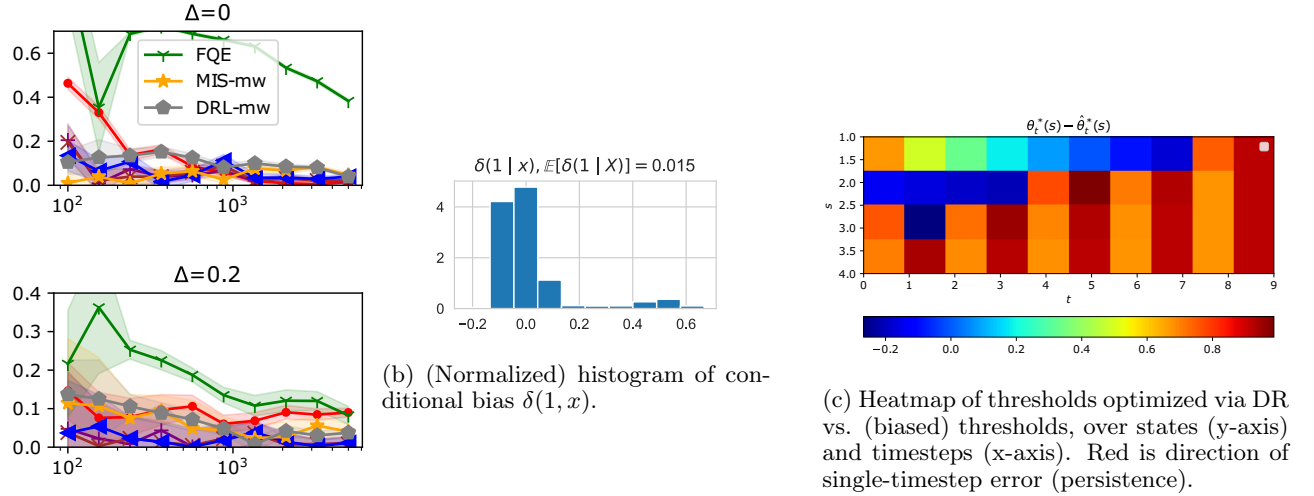


Figure 2: Policy evaluation and optimization as more trajectories (x-axis) are collected of T -horizon selling in contextual and capacitated dynamic pricing, example 1; specification of eq. (10).



(a) $\Delta = 0.2$; further OPE comparison on a different DGP.

Figure 3: Conditions of Theorem 3 for error persistence.

consider a (time and state-stationary) evaluation policy $\pi_e(1 | x) = \sigma(0.25(\beta^\top x))$. The time horizon is 10 timesteps, with initial state capacity $s_0 = 4$.

Policy evaluation and optimization. In Figure 2 we generate different outcome models with increasing levels of misspecification Δ , evaluate $V_0^{\pi_e}(s_0)$ by Monte Carlo rollouts with $N = 10000$ trajectories. We compare DM with logistic regression nuisance, DR doubly-robust with logistic regression, DR-nonpara with nonparametric nuisance, and IPW, inverse propensity weighting. (See Section D for further comparison including other baselines and policy optimization in the well-specified $\Delta = 0$ case, where the variance drawbacks of IPW do worse than model-based approaches.) Figure 2a considers off policy evaluation, with absolute relative error on the y-

axis and trajectory size on the x-axis (log grid from $N = 50, \dots, 5000$). When $\Delta = 0.2$ the logistic outcome model is misspecified, but orthogonality and the well-specified propensity score ensures estimates are asymptotically unbiased. Similar to other DR settings, although incorporating the outcome model reduces variance, incorporating a misspecified outcome model does worse than just using well-specified IPW, but we see faster convergence from the flexible, nonparametric nuisance which outperforms well-specified IPW. We also compared to nonparametric baselines FQE Le et al. (2019), and modified stateful versions of MIS Yin et al. (2021) and DRL Kallus and Uehara (2019a). However, in this simple setting, the highly flexible nuisance estimators overfit and fail (incurring 40-50% absolute error). We discuss these baselines in greater detail in “OPE comparison” in a more favor-

able data-generating process.

We then consider policy optimization in Figure 2b with a rich policy class to avoid misspecification error issues. Motivated by eq. (8), observe that the optimal threshold policy on the true $\Delta\mu$ is an affine transform relative to a threshold on the estimated $\Delta\hat{\mu}$ (possibly misspecified, hence biased), with an x -conditional term for the conditional error. We approximate optimizing over policies $\mathbb{I}[\Delta\mu > \theta]$ by ranging over all thresholds on $\mathbb{I}[\Delta\hat{\mu} > \theta']$; this approximates the x -conditional error term of eq. (8) by a constant. This is similar to a contextual version of “bid-price” policies [Gallego et al. (2019)]. We optimize over the class of threshold policies on $\Delta\hat{\mu}$ by enumerating thresholds and evaluating via the estimate from Proposition 2, so the functional specification does depend on the (unadjusted) nuisance estimation. (Therefore the VC dimension of this depends on the VC dimension of the outcome nuisance). The y -axis depicts out of sample value (higher is better) averaged over 48 replications. Both DR and IPW (inverse propensity-weighted) estimates translate to improvements in optimized policy value. We see dramatic benefits of DR when $\Delta = 0.2$. For small dataset sizes, IPW suffers from high variance as expected. Therefore, DR and its variance reduction estimates achieve sizeable improvements for small amounts of data. As the amount of data grows larger, the performance of IPW nears that of DR asymptotically. The DM plug-in approach remains biased and achieves worse performance, even asymptotically.

OPE comparison. We compare to state-of-the-art OPE: FQE of [Le et al. (2019)] which does not use the “stateful” structure, and we also derive “strong baselines” that leverage some of the structure (MIS-mw [Yin et al. (2021)], DRL-mw [Kallus and Uehara (2019b)]). (We reiterate our core contribution is not in general off-policy evaluation but in deriving improvements for this specific structure). For example, observe that since x is exogenously generated the finite-horizon *state-action density ratio* is *independent of x* . We endow MIS-mw and DRL-mw with this structural information (see appendix Section D for details). As the general OPE literature prescribes, we use nonparametric nuisances, e.g. multi-layer perceptron with scikit-learn defaults for all nuisance predictors. We consider a more favorable DGP for OPE comparison in Figure 3a, using eq. (10) with $p = 5$ and $\beta = [-0.53, -0.56, -0.10, 0.40, 0.74]$, $\beta_0 = -2.39$. MIS-mw does well overall, although empirically we find other DGPs where MIS-mw underperforms FQE. In the misspecified setting, our doubly robust estimators outperform MIS-mw. FQE, which fits next time-step $Q(s, x, a)$, appears to converge but much slower than our approaches. The gap between FQE and DM for

the (slightly misspecified) case precisely illustrates the benefits of encoding problem structure in Equation (1).

Assessing the structural conditions of Theorem 3 in practice. In Figure 3 we investigate assumptions made in Example 4 (e.g. uniformity of error direction $\delta(1, x)$) that do not hold exactly. Figure 3b plots δ : although it is symmetrically distributed for most x , there is overall marginal error in the expected direction. In the empirical example we optimize over marginal thresholds and so we expect, marginalizing over x , the directional error condition is satisfied. In Figure 3c, we show a heatmap of $\theta_t^*(s) - \hat{\theta}_t(s)$ over timesteps and state values. As the analysis suggests, for $s > 2$ for most timesteps the error persists: red indicates regions where naive thresholds are in the same direction, relative to the optimal threshold, and hence the error persists rather than attenuates over time.

7 Concluding remarks.

By studying the causal structure of practically relevant problems in operations, we developed specialized off policy evaluation and optimization which demonstrate the offline version of such problems is easier than a generic MDP. We show analytically that confounding matters, and verify our approach, reducing from $T + 1$ nuisances to 2 and estimating the expectation of a transition via the expectation over a population, achieves practical benefits.

Violations of assumptions and extensions Consider a specific violation of Asn. 1 with *state-dependent contexts*, where X_t may depend on S_t but nothing else; if Asn. 2 additionally holds, system state S is not a confounder, but covariate shifts in S induce shifts in contexts X that are confounders. Therefore eq. 3 is a biased estimate of the transition probability at a specific state due to covariate shift. We model this bias as single-timestep OPE under *covariate shift in contexts*, since the only confounding bias is due to integrating over the observational state-context occupancy distribution. The density ratio factorizes so the only unknown is $\frac{d_t^c(S)}{d_t^b(S)}$. This suggests we may use ideas based on robustness: under some additional restrictions on how far $\frac{d_t^c(S)}{d_t^b(S)}$ can vary from 1, we can use the optimization schemes suggested in Footnote 1 for adversarial arrivals. We can optimize for robust transitions in the *Marginal MDP*, conduct a robust backwards induction, and return a covariate-shift robust policy.

Acknowledgments This work was supported by the National Science Foundation under Grant No. 1846210.

References

- S. Agrawal and R. Jia. Learning in structured mdps with convex cost functions: Improved regret bounds for inventory management. In *Proceedings of the 2019 ACM Conference on Economics and Computation*, pages 743–744, 2019.
- S. Agrawal, N. R. Devanur, and L. Li. An efficient algorithm for contextual bandits with knapsacks, and an extension to concave objectives. In *Conference on Learning Theory*, pages 4–18, 2016.
- A. Antos, R. Munos, and C. Szepesvari. Fitted q-iteration in continuous action-space mdps. In *Neural Information Processing Systems*, 2007.
- A. Badanidiyuru, R. Kleinberg, and A. Slivkins. Bandits with knapsacks. *Journal of the ACM (JACM)*, 65(3):1–55, 2018.
- G.-Y. Ban and N. B. Keskin. Personalized dynamic pricing with machine learning: High dimensional features and heterogeneous elasticity. *Forthcoming, Management Science*, 2020.
- D. P. Bertsekas and J. N. Tsitsiklis. *Neuro-dynamic programming*. Athena Scientific, 1996.
- O. Besbes and A. Zeevi. Blind network revenue management. *Operations research*, 60(6):1537–1550, 2012.
- A. F. Bibaut, I. Malenica, N. Vlassis, and M. J. van der Laan. More efficient off-policy evaluation through regularized targeted learning. *arXiv preprint arXiv:1912.06292*, 2019.
- K. Bimpikis, O. Candogan, and D. Saban. Spatial pricing in ride-sharing networks. *Operations Research*, 67(3):744–769, 2019.
- H. Bojun. Steady state analysis of episodic reinforcement learning. *arXiv preprint arXiv:2011.06631*, 2020.
- R. L. Bray. The multisecretary problem with continuous valuations. *arXiv preprint arXiv:1912.08917*, 2019.
- G. Brockman, V. Cheung, L. Pettersson, J. Schneider, J. Schulman, J. Tang, and W. Zaremba. Openai gym. *arXiv preprint arXiv:1606.01540*, 2016.
- P. Bumpensanti and H. Wang. A re-solving heuristic with uniformly bounded loss for network revenue management. *Management Science*, 2020.
- J. Chen and N. Jiang. Information-theoretic considerations in batch reinforcement learning. In *International Conference on Machine Learning*, pages 1042–1051. PMLR, 2019.
- X. Chen, Z. Owen, C. Pixton, and D. Simchi-Levi. A statistical learning approach to personalization in revenue management. *Management Science*, 2021.
- V. Chernozhukov, D. Chetverikov, M. Demirer, E. Duflo, C. Hansen, W. Newey, and J. Robins. Double/debiased machine learning for treatment and structural parameters, 2018.
- M. C. Cohen, I. Lobel, and R. P. Leme. Feature-based dynamic pricing. *EC*, 2016(10.1145):2940716–2940728, 2016.
- Y. Duan, Z. Jia, and M. Wang. Minimax-optimal off-policy evaluation with linear function approximation. In *International Conference on Machine Learning*, pages 2701–2709. PMLR, 2020.
- M. Dudik, D. Erhan, J. Langford, and L. Li. Doubly robust policy evaluation and optimization. *Statistical Science*, 2014.
- I. El Shar and D. Jiang. Lookahead-bounded q-learning. In *International Conference on Machine Learning*, pages 8665–8675. PMLR, 2020.
- Y. Emek, R. Lavi, R. Niazadeh, and Y. Shi. Stateful posted pricing with vanishing regret via dynamic deterministic markov decision processes. *Advances in Neural Information Processing Systems*, 33, 2020.
- A.-m. Farahmand and C. Szepesvári. Regularized least-squares regression: Learning from a β -mixing sequence. *Journal of Statistical Planning and Inference*, 142(2):493–505, 2012.
- G. Gallego and G. Van Ryzin. A multiproduct dynamic pricing problem and its applications to network yield management. *Operations research*, 45(1):24–41, 1997.
- G. Gallego, H. Topaloglu, et al. *Revenue management and pricing analytics*, volume 209. Springer, 2019.
- L. Györfi, M. Kohler, A. Krzyżak, and H. Walk. *A distribution-free theory of nonparametric regression*, volume 1. Springer, 2002.
- Y. Hu, N. Kallus, and M. Uehara. Fast rates for the regret of offline reinforcement learning. In *Conference on Learning Theory*, 2021.
- W. T. Huh, R. Levi, P. Rusmevichientong, and J. B. Orlin. Adaptive data-driven inventory control with censored demand based on kaplan-meier estimator. *Operations Research*, 59(4):929–941, 2011.
- A. Javanmard and H. Nazerzadeh. Dynamic pricing in high-dimensions. *arXiv preprint arXiv:1609.07574*, 2016.
- N. Jiang and L. Li. Doubly robust off-policy value evaluation for reinforcement learning. *Proceedings of the 33rd International Conference on Machine Learning*, 2016.
- C. Jin, Z. Allen-Zhu, S. Bubeck, and M. I. Jordan. Is q-learning provably efficient? *arXiv preprint arXiv:1807.03765*, 2018.

- N. Kallus and M. Uehara. Double reinforcement learning for efficient off-policy evaluation in markov decision processes. *arXiv preprint arXiv:1908.08526*, 2019a.
- N. Kallus and M. Uehara. Efficiently breaking the curse of horizon: Double reinforcement learning in infinite-horizon processes. *arXiv preprint arXiv:1909.05850*, 2019b.
- N. Kallus and A. Zhou. Minimax-optimal policy learning under unobserved confounding. *Management Science*, 2020.
- E. H. Kennedy, S. Balakrishnan, and L. Wasserman. Semiparametric counterfactual density estimation. *arXiv preprint arXiv:2102.12034*, 2021.
- T. Kitagawa and A. Tetenov. Empirical welfare maximization. 2015.
- H. Le, C. Voloshin, and Y. Yue. Batch policy learning under constraints. In *International Conference on Machine Learning*, pages 3703–3712. PMLR, 2019.
- S. Meyn and S. P. Meyn. *Control techniques for complex networks*. Cambridge University Press, 2008.
- S. P. Meyn and R. L. Tweedie. *Markov chains and stochastic stability*. Springer Science & Business Media, 2012.
- M. Mohri, A. Rostamizadeh, and A. Talwalkar. *Foundations of machine learning*. MIT press, 2018.
- R. Munos. Error bounds for approximate policy iteration. In *ICML*, volume 3, pages 560–567, 2003.
- R. Munos. Performance bounds in l_p -norm for approximate value iteration. *SIAM journal on control and optimization*, 46(2):541–561, 2007.
- R. Munos and C. Szepesvári. Finite-time bounds for fitted value iteration. *Journal of Machine Learning Research*, 9(5), 2008.
- X. Nie, E. Brunskill, and S. Wager. Learning when-to-treat policies. *Journal of the American Statistical Association*, pages 1–18, 2020.
- F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, et al. Scikit-learn: Machine learning in python. *the Journal of machine Learning research*, 12:2825–2830, 2011.
- D. Pollard. *Empirical processes: theory and applications*. NSF-CBMS regional conference series in probability and statistics, 1990.
- W. B. Powell. *Approximate Dynamic Programming: Solving the curses of dimensionality*, volume 703. John Wiley & Sons, 2007.
- S. Qiang and M. Bayati. Dynamic pricing with demand covariates. *Available at SSRN 2765257*, 2016.
- B. Scherrer. Approximate policy iteration schemes: a comparison. In *International Conference on Machine Learning*, pages 1314–1322. PMLR, 2014.
- V. Shah, J. Blanchet, and R. Johari. Semi-parametric dynamic contextual pricing. *arXiv preprint arXiv:1901.02045*, 2019.
- A. Swaminathan and T. Joachims. The self-normalized estimator for counterfactual learning. *Proceedings of NIPS*, 2015.
- Z. Tang, Y. Feng, L. Li, D. Zhou, and Q. Liu. Doubly robust bias reduction in infinite horizon off-policy estimation. *arXiv preprint arXiv:1910.07186*, 2019.
- P. Thomas and E. Brunskill. Data-efficient off-policy policy evaluation for reinforcement learning. In *International Conference on Machine Learning*, pages 2139–2148, 2016.
- A. W. Van Der Vaart and J. A. Wellner. Weak convergence. In *Weak convergence and empirical processes*, pages 16–28. Springer, 1996.
- R. Vershynin. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge University Press, 2018.
- S. Wager and S. Athey. Efficient policy learning. 2017.
- M. J. Wainwright. *High-dimensional statistics: A non-asymptotic viewpoint*, volume 48. Cambridge University Press, 2019.
- T. Xie, Y. Ma, and Y.-X. Wang. Towards optimal off-policy evaluation for reinforcement learning with marginalized importance sampling. *arXiv preprint arXiv:1906.03393*, 2019.
- Y. Yang and A. Barron. Information-theoretic determination of minimax rates of convergence. *Annals of Statistics*, pages 1564–1599, 1999.
- M. Yin, Y. Bai, and Y.-X. Wang. Near-optimal provable uniform convergence in offline policy evaluation for reinforcement learning. In *International Conference on Artificial Intelligence and Statistics*, pages 1567–1575. PMLR, 2021.
- B. Zhang, A. A. Tsiatis, E. B. Laber, and M. Davidian. Robust estimation of optimal dynamic treatment regimes for sequential treatment decisions. *Biometrika*, 100(3):681–694, 2013.
- Y.-Q. Zhao, D. Zeng, E. B. Laber, and M. R. Kosorok. New statistical learning methods for estimating optimal dynamic treatment regimes. *Journal of the American Statistical Association*, 110(510):583–598, 2015.
- Z. Zhou, S. Athey, and S. Wager. Offline multi-action policy learning: Generalization and optimization. *arXiv preprint arXiv:1810.04778*, 2018.

Outline of appendix

Section [A](#) analyzes the marginal MDP construction of Section [4](#)

Section [B](#) provides analysis of the main estimation strategy and policy learning sample complexity analysis, e.g. proofs of Section [5](#).

Section [C](#) provides additional discussion.

Section [D](#) includes additional empirics and discussion of computational details.

A Proofs for Section [4](#)

Proof of Proposition [1](#). We argue via backwards induction, though the proof follows largely from construction of the marginal MDP. Consider the last timestep, $t = T$: equivalence follows by construction. Next consider the inductive step. The inductive hypothesis posits equivalence of policy values $V_{t+1}(s)$ in the marginal MDP and $\hat{V}_{t+1}(s)$, the marginalized policy value in the full-information MDP, as well as equivalence of optimal policies so that $V_{t+1}(s) = \hat{V}_{t+1}(s)$.

Recall that each policy $\pi \in \Pi$ in $\mathcal{M}$ corresponds to an action $\pi \in \Pi$ of $\widetilde{\mathcal{M}}$ with transition probability:

$$\tilde{P}(s' | s, \pi) = \int \int P(s' | s, a, x) \pi(a | x) f(x) dx \quad (11)$$

Equivalence of the policy values follows from backwards induction of expected rewards and definition of the marginal MDPs. Equivalence of optimal policies follows from equivalence of first moments and equivalence of the policy classes. Equivalence of the inductive step follows from equivalence of policy classes, identically distributed arrivals X_i within a timestep. $\square$

Proof of Proposition [2](#). Double robustness is standard and follows standard arguments in the single time-step policy learning literature [Dudik et al. \(2014\)](#); [Wager and Athey \(2017\)](#). The claim follows by observing that by the restricted causal structure Assumption [2](#), $Y(a) \perp\!\!\!\perp S | X$, so it is sufficient to adjust for IPW weights $P(A = a | X = x)$.

For completeness, we verify double-robust unbiasedness properties of $\hat{V}_t(s)$. The proof follows by backward induction. The fact that $V_{T+1}(s) = 0$ and estimation is unbiased follows from double-robustness in the single-timestep setting for $\hat{V}_T(s)$ for all $s \in \mathcal{S}$.

Next we show the inductive step. Suppose $\hat{V}_{t+1}(s)$ is unbiased, for all $s \in \mathcal{S}$. We require that the evaluation error is independent ($\hat{V}_{t+1}^{\text{DR}}(s') - V_{t+1}(s')$) from the nuisance evaluation error, which can be easily satisfied by appropriate sample splitting.

If μ is unbiased:

$$\mathbb{E}[\hat{V}_t(s)] = \mathbb{E} \left[\sum_{s', y | s} \sum_a \pi_e(a | x) \mu_t(y | a, x) (R(y) + \hat{V}_{t+1}^{\text{DR}}(s')) \right] \quad (12)$$

$$+ \mathbb{E} \left[\sum_{s', y | s} \sum_a \pi_e(a | x) \left(\frac{\mathbb{I}[A_t = a]}{\pi_b(a | x)} (\mathbb{I}[Y = y] - \mu_t(y | a, x)) \right) (R(y) + V_{t+1}(s')) \right] \quad (13)$$

$$+ \mathbb{E} \left[\sum_{s', y | s} \sum_a \pi_e(a | x) \left(\frac{\mathbb{I}[A_t = a]}{\pi_b(a | x)} (\mathbb{I}[Y = y] - \mu_t(y | a, x)) \right) (\hat{V}_{t+1}^{\text{DR}}(s') - V_{t+1}(s')) \right] \quad (14)$$

Note eq. [\(12\)](#) = $V_t(s)$ by well-specification, eq. [\(13\)](#) = 0 by unbiasedness of μ_t and cross-fitting so the estimation errors are independent, and the first term of eq. [\(14\)](#) is expectation-0 by the previous argument and $\mathbb{E}[(\hat{V}_{t+1}^{\text{DR}}(s') - V_{t+1}(s'))] = 0$ by the induction hypothesis and cross-fitting.

If e is unbiased:

$$\begin{aligned}
 & \mathbb{E}[\widehat{V}_t(s)] \\
 &= \mathbb{E} \left[\sum_{s', y|s} \sum_a \pi_e(a | x) \left(\underbrace{\frac{\mathbb{I}[A_t = a]}{\pi_b(a | x)} \mathbb{I}[Y = y]}_{\text{unbiased IPW estimator}} + \underbrace{\mu_t(y | a, x) \left(1 - \frac{\mathbb{I}[A_t = a]}{\pi_b(a | x)} \right)}_{=0 \text{ by unbiasedness of } e} \right) (R(y) + V_{t+1}(s')) \right] \\
 &+ \mathbb{E} \left[\sum_{s', y|s} \sum_a \pi_e(a | x) \left(\underbrace{\mu_t(y | a, x) \left(1 - \frac{\mathbb{I}[A_t = a]}{\pi_b(a | x)} \right)}_{=0 \text{ by unbiasedness of } e} \right) \underbrace{(\widehat{V}_{t+1}^{\text{tDR}}(s') - V_{t+1}(s'))}_{=0 \text{ by inductive hypothesis}} \right]
 \end{aligned}$$

□

B Proofs for Section 5

B.1 Sample complexity proofs

B.1.1 Concentration preliminaries

We introduce the uniform convergence setup we use to provide tail inequalities. We will apply a standard chaining argument with Orlicz norms and introduce some notations from standard references, e.g. [Vershynin \(2018\)](#); [Pollard \(1990\)](#); [Wainwright \(2019\)](#). A function $\phi : [0, \infty) \rightarrow [0, \infty)$ is an Orlicz function if ϕ is convex, increasing, and satisfies $\phi(0) = 0, \phi(x) \rightarrow \infty$ as $x \rightarrow \infty$. For a given Orlicz function ϕ , the Orlicz norm of a random variable X is defined as $\|X\|_\phi = \inf\{t > 0: \mathbb{E}[\Phi(\|X\| | t)] \leq 1\}$. The Orlicz norm $\|Z\|_\Phi$ of random variable Z is defined by $\|Z\|_\Phi = \inf\{C > 0: \mathbb{E}[\Phi(Z/C)] \leq 1\}$. A constant bound on $\|Z\|_\Phi$ constrains the rate of decrease for the tail probabilities by the inequality $\mathbb{P}(|Z| \geq t) \leq 1/\Phi(t/C)$ if $C = \|Z\|_\Phi$. For example, choosing the Orlicz function $\Phi(t) = \frac{1}{5} \exp(t^2)$ results in bounds by subgaussian tails decreasing like $\exp(-Ct^2)$, for some constant C .

We next introduce the tail inequalities that use a standard chaining argument to control uniform convergence over $\pi \in \Pi$. Let n denote a generic dataset size (we will later on apply the results with $n = NT$.) The data are $(X_{1:n}, A_{1:n}, Y_{1:n})$ and $f_i(\pi)$ is a function of (X_i, A_i, Y_i) . Define the function class $\mathcal{F}(X_{1:n}, A_{1:n}, Y_{1:n}) = \{(f_i(\pi), \dots, f_n(\pi)) : \pi \in \Pi\}$.

For this section, we consider maximal inequalities for the function classes for the enveloped policy class $\mathcal{F}$. Let $\epsilon_i \in \{-1, +1\}$, be iid Rademacher variables (symmetric Bernoulli random variables with value $-1, +1$ with probability $\frac{1}{2}$), distributed independently of all else. We use the following application of chaining with a bounded envelope function, due to [Pollard, 1990](#), Eqn. 7.3). (Using different measures of functional complexity for multi-class predictors, such as Natarajan dimension, simply changes the constants in the final bound.)

Theorem A (Uniform convergence of policy function π over envelope class $\mathcal{F}$). *Let $f(\pi) \leq \|F\|_2 \leq C$ be a bound on the envelope function for $f \in \mathcal{F}$. Then for n large enough, where d_{vc} is the VC-dimension,*

$$\sup_{f \in \mathcal{F}} \left| \frac{1}{n} \sum_{i=1}^n (f_i(\pi) - \mathbb{E}[f(\pi)]) \right| \leq 9/2 C \sqrt{\frac{d_{\text{vc}} \log(5/\delta)}{n}} \quad (15)$$

The next variant is a modification of Thm. [A](#) which only uses only moment bounds for the envelope function at the expense of weaker controls of the tails of the supremum process.

Theorem B (Uniform convergence with L_p norm of envelope function.). *For an absolute constant $C_{d_{\text{vc}}}$ which depends only on d_{vc} , where F_n is the envelope for $f \in \mathcal{F}$ and d_{vc} is the VC dimension, $\mathbb{E}[\sup_{f \in \mathcal{F}} |\frac{1}{n} \sum_{i=1}^n (f_i(\pi) - \mathbb{E}[f(\pi)])|] \leq C'_{d_{\text{vc}}} \sqrt{d_{\text{vc}}} - \mathbb{1}\mathbb{E}[F_n]$.*

We also state a standard lemma used for sample splitting, as appears in [Chernozhukov et al. \(2018\)](#), without proof.

Lemma 1 (Conditional convergence implies unconditional). *Let $\{X_m\}$ and $\{Y_m\}$ be sequences of random vectors.*

- *If for $\epsilon_m \rightarrow 0$, $\mathbb{P}(\|X_m\| > \epsilon_m \mid Y_m) \rightarrow_p 0$, then $\mathbb{P}(\|X_m\| > \epsilon_m) \rightarrow 0$. This occurs if $\mathbb{E}[\|X_m\|^q / \epsilon_m^q \mid Y_m] \rightarrow_p 0$ for some $q \geq 1$ by Markov's inequality.*
- *Let $\{A_m\}$ be a sequence of positive constants. If $\|X_m\| = O_p(A_m)$ conditional on Y_m , that for any $\ell_m \rightarrow \infty$, $\mathbb{P}(\|X_m\| > \ell_m A_m \mid Y_m) \rightarrow_p 0$, then $\|X_m\| = O_p(A_m)$ unconditionally, namely, that for any $\ell_m \rightarrow \infty$, then $\mathbb{P}(\|X_m\| > \ell_m A_m) \rightarrow 0$.*

Proof of Thm. A. We first bound the deviations uniformly over the policy class and introduce the following empirical processes,

$$M = \sup_{f \in \mathcal{F}} \left| \sum_{i=1}^n (f_i(\pi) - \mathbb{E}[f(\pi)]) \right|, \quad L = \sup_{f \in \mathcal{F}} \left| \sum_{i=1}^n \epsilon_i f_i(\pi) \right|.$$

By a standard symmetrization argument, applying Jensen's inequality for the convex function Φ of the symmetrized process (e.g. Theorem 2.2 of Pollard (1990)), we may bound the Orlicz norm of the maxima of the empirical process by the symmetrized process, conditional on the observed data: $\mathbb{E}[\Phi(M)] \leq \mathbb{E}[\Phi(2L)]$. Taking Orlicz norms with $\Phi(t) = \frac{1}{5} \exp(t^2)$, we apply a tail inequality on the Orlicz norm of the symmetrized process $\Phi(2L)$, under the assumption of bounded outcomes. Applying Dudley's inequality to the symmetrized empirical process L , (e.g. Theorem 3.5 of Pollard (1990)), we have that

$$\mathbb{E}_\epsilon [\exp(L^2/J^2) \mid \mathcal{D}] \leq 5 \quad \text{for} \quad J = 9 \|F\|_2 \int_0^1 \sqrt{\log(D(\|F\|_2 \zeta, \mathcal{F}(X_{1:n})))} d\zeta. \quad (16)$$

By Markov's inequality, we have that $\mathbb{P}(\frac{1}{n}L > t) \leq 5 \exp(-t^2 n^2 / \|L\|_2^2) = 5 \exp(-t^2 n / J^2 C^2)$, so that therefore, bounding the Dudley entropy integral by the VC dimension via Pollard (1990) eqn. 7.8),

$$\frac{1}{n} M \leq \frac{9/2 C \sqrt{d_{\text{vc}} \log(5/\delta)}}{\sqrt{n}}.$$

□

Proof of Thm. B. By standard results on covering numbers and VC dimension, e.g. Van Der Vaart and Wellner (1996); Wainwright (2019), for a VC-class of functions with measurable envelope function F and $p \geq 1$, for any probability measure Q with $\|F\|_{Q,p} > 0$, $N(\epsilon \|F\|_{Q,p}, \mathcal{F}, L_p(Q)) \leq A(d_{\text{vc}})(\frac{1}{\epsilon})^{p(d_{\text{vc}}-1)}$, where $C_{d_{\text{vc}}}$ is an absolute constant that depends only on the vc-dimension. By Pollard (1990) eqn. 7.8),

$$\begin{aligned} \mathbb{E}[\sup_{\pi} \left| \frac{1}{n} \sum_{i=1}^n (f_i(\pi) - \mathbb{E}[f(\pi)]) \right|^p] &\leq \left(18 C_p \int_0^1 \sqrt{\log(C_{d_{\text{vc}}} (1/\epsilon)^{p(d_{\text{vc}}-1)})} dx \right) \mathbb{E}[\|F_n\|^p] \\ &\leq C'_{d_{\text{vc}}} \sqrt{d_{\text{vc}}} - 1 \mathbb{E}[\|F_n\|^p] \end{aligned}$$

The claim follows by taking $p = 2$.

□

Preliminaries Having modeled the marginal MDP, the analysis proceeds via a standard performance difference lemma which provides an additive decomposition of the regret for finite horizons; and single-timestep causal inference. For example, this appears in Jin et al. (2018); we simply include the full statement for completeness and verify for our setting. The ease of analysis is due to our reduction to the marginal MDP: the original MDP may have a continuous state so that tabular MDP analysis is not possible. Recall that we denote $\tilde{P}$ as the transition matrices and estimated transition matrices corresponding to the marginal MDP (e.g. transitions between system states). In this section, for brevity we let $\hat{P}$ denote the empirical counterpart of $\tilde{P}$, e.g. estimating eq. (2) via some estimation strategy (IPW weighting, doubly robust, or plug-in estimation; typically we focus on the doubly-robust estimator). Similarly, $\hat{M}$ denotes the empirical MDP model with $\hat{P}$. We introduce notation for indexing into entries after evaluating the transition operator,

$$(\tilde{P}V)(s, \pi) = \mathbb{E}_{s' \sim \tilde{P}(\cdot | s, \pi)} V(s') = \mathbb{E}[\tilde{R}(s, \pi)] + \sum_{(s', y) | s} P(Y = y \mid \pi)(V(s'))$$

We also introduce notation for indexing the difference between applying the true transition operator and the empirical estimate thereof, for any *generic* $|\mathcal{S}|$ -vector v and policy π :

$$((\tilde{P} - \hat{P})v)(s, \pi) = \sum_{(s', y) | s} \left(P(Y = y | \pi) - \hat{P}(Y = y | \pi) \right) v(s').$$

Lemma 2 (Additive decomposition of finite-horizon value iteration). *For any policies $\tilde{\pi}, \pi$ and any $(s, t') \in \mathcal{S} \times [T]$,*

$$\begin{aligned} & \tilde{V}_{t'}^{\tilde{\pi}}(s) - \hat{V}_{t'}^{\pi}(s) \\ &= \mathbb{E}_{\tilde{M}, \tilde{\pi}} \left[\sum_{t=t'}^T \mathbb{E}[\tilde{R}(s, \pi_t)] - \hat{\mathbb{E}}[\tilde{R}(s, \pi_t)] + \sum_{(s', y) | s} \left(P(Y = y | \pi_t) - \hat{P}(Y = y | \pi_t) \right) \tilde{V}_{t+1}^{\pi}(s') \mid s_t = s \right] \end{aligned}$$

B.2 Sample complexity analysis

We first provide a generalization bound in an *oracle nuisance* case where we use the true conditional expectations μ, e rather than estimated counterparts $\hat{\mu}, \hat{e}$.

$$\begin{aligned} \Gamma_t^{i,*}(y | a) &:= \frac{\mathbb{I}[Y_t^i = y] - \mu(y | A_t^i, X_t^i)}{e(A_t^i | X_t^i)} \mathbb{I}[A_t^i = a] + \mu(y | a, X_t^i) \\ \hat{P}(Y = y | \pi) &:= \frac{1}{NT} \sum_{t=1}^T \sum_{i=1}^N \sum_{a \in \mathcal{A}} \pi(a | X_t^i) \hat{\Gamma}_t^{i,*}(y | a) \end{aligned} \quad (17)$$

Theorem 4 (Sample complexity and rate double-robustness for oracle estimator). *Suppose $e(a | x) \leq \nu^{-1}$ uniformly over a, x (overlap).*

Then w.p. $\geq 1 - \delta$, for $\tilde{V}_0^{\hat{\pi}^}$ optimized via Algorithm 1 with oracle nuisance estimator eq. (17),*

$$\tilde{V}_0^{\pi^*} - \tilde{V}_0^{\hat{\pi}^*} \leq \frac{\nu d_{\text{vc}} R^{\max}(T + 1/2 T^2)^{9/2} \sqrt{\log(5T|\mathcal{Y}|/\delta)}}{\sqrt{NT}}$$

Proof of Theorem 4. In the analysis, we leverage single-timestep uniform convergence arguments. By Assumption 1, X is drawn exogenously/independently of all else, so the X data is iid. We model the dataset as drawn from multiple behavior policies, so that at timestep t , $S \sim \rho_t^{\pi_b}(S)$, $X \sim f_t$, $A \sim \pi_b(s, x)$, and $Y | A, X$ is drawn from the contextual response model. Therefore the data tuples $(S_{i,t}, X_{i,t}, A_{i,t}, Y_{i,t})$ are viewed as independent draws from this process. Since $Y | A, X \perp\!\!\!\perp S_t$ and therefore we only need to control for X_t such that $e(X_t)$ is not a function of S , the functions we use in our estimator, defined with respect only to the $(X_{i,t}, A_{i,t}, Y_{i,t})$ data, admit analysis via iid/single-stage empirical process techniques.

We define the following function classes conditional on all the data, $(X_{1:NT}, A_{1:NT}, Y_{1:NT})$. For π , we consider a shifted function class with an envelope function: let $f_{it}(\pi) = \pi(A_t^i | X_t^i) \mathbb{I}[A_t^i = a] \mathbb{I}[Y_t^i = y]$ where

$$\mathcal{F}(X_{1:NT}, A_{1:NT}, Y_{1:NT}) = \{(f_1^1(\pi), \dots, f_t^i(\pi), \dots, f_T^N(\pi)) : \pi \in \Pi\}.$$

Step 1: Error decomposition:

We decompose the error. By optimality of $\hat{\pi}^*$, $\tilde{V}_0^{\pi^*} - \hat{V}_0^{\hat{\pi}^*} \leq 0$ and the triangle inequality:

$$\begin{aligned} \tilde{V}_0^{\pi^*} - \tilde{V}_0^{\hat{\pi}^*} &= \tilde{V}_0^{\pi^*} - \hat{V}_0^{\pi^*} + \hat{V}_0^{\pi^*} - \hat{V}_0^{\hat{\pi}^*} + \hat{V}_0^{\hat{\pi}^*} - \tilde{V}_0^{\hat{\pi}^*} \\ &\leq \left| \tilde{V}_0^{\pi^*} - \hat{V}_0^{\pi^*} \right| + \left| \hat{V}_0^{\hat{\pi}^*} - \tilde{V}_0^{\hat{\pi}^*} \right| \end{aligned} \quad (18)$$

Step 2: Uniform convergence over π and $\tilde{V}^{\pi}$:

We apply the additive decomposition of Lemma 2 and obtain a uniform bound on $\sup_{\pi \in \Pi} \tilde{V}_0^{\tilde{\pi}}(s) - \hat{V}_0^{\pi}(s)$, which we apply twice to the terms of Equation (18).

For π , we consider a shifted function class with an envelope function: define

$$f_i(\pi) = \pi(A_i | X_i) \mathbb{I}[A_i = a] \mathbb{I}[Y_i = y] e^{-1}(a | X_i)$$

$$\mathcal{F}(X_{1:NT}, A_{1:NT}, Y_{1:NT}) = \{(f_1(\pi), \dots, f_{NT}(\pi)) : \pi \in \Pi\}.$$

By Lemma 2, for any policy π , consider the additive decomposition:

$$\sup_{\pi \in \Pi} \tilde{V}_0^{\tilde{\pi}}(s) - \hat{V}_0^{\pi}(s) = \sup_{\tilde{\pi} \in \Pi} \mathbb{E}_{\hat{M}, \tilde{\pi}} \left[\sum_{t=0}^T (\tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t)) + (\tilde{P}_t - \hat{P}_t)(\tilde{V}_{t+1}^{\pi}(s_t, \pi_t) | s_t = s) \right]$$

Note that the dependence on the state distribution is not relevant (that is, the expectation under $(\hat{M}, \hat{\pi})$ since the estimator uses the same data at every state s , and hence it suffices to consider uniformity over $\mathcal{Y}$ and equivalently evaluate the supremum over corresponding transitions.

$$\begin{aligned} & \sup_{(\pi, \pi') \in \Pi \times \Pi} \sum_{t=0}^T \sup_{y \in \mathcal{Y}} \left\{ (\tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t)) + (P(Y = y | \pi_t) - \hat{P}(Y = y | \pi_t))(\tilde{V}_{t+1}^{\pi'}(s_t, \pi_t)) \right\} \\ & \leq \sum_{t=0}^T \sup_{(\pi_t, \pi_{t+1:T}^{(t)}) \in \Pi_t \times \Pi_{t+1:T}} \sup_{y \in \mathcal{Y}} \left\{ (\tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t)) + (P(Y = y | \pi_t) - \hat{P}(Y = y | \pi_t))(\tilde{V}_{t+1}^{\pi_{t+1:T}^{(t)}}(s_t, \pi_t)) \right\} \end{aligned} \quad (19)$$

$$\leq \sum_{t=0}^T \sup_{y \in \mathcal{Y}} \left\{ \sup_{\pi_t \in \Pi_t} |(\tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t))| + \sup_{\pi_t \in \Pi_t} |(T - t - 1) \cdot (P(Y = y | \pi_t) - \hat{P}(Y = y | \pi_t))| \right\} \quad (20)$$

Since forward transitions marginalize to $\tilde{V}_{t+1}(s')$, in the last line we observed that taking the supremum over $\pi_{t+1:T}$ can only enlarge the fixed envelope function, $\|F\|_2$, but does not actually affect the empirical process analysis. Under assumption of a product set of policies across states, uniform convergence under $\pi_t \in \Pi_t$ also establishes uniform convergence for state-dependent policies.

Step 3: applying the concentration inequalities. We bound each of the above terms by Thm. A, applying a high probability bound with $\delta' = \frac{\delta}{T}$, and finally take a union bound over $\mathcal{Y}$ and each summand, in order to obtain the following bound which holds with probability $> 1 - \delta$,

$$\begin{aligned} \sup_{\tilde{\pi} \in \Pi} \tilde{V}_0^{\tilde{\pi}}(s) - \hat{V}_0^{\pi}(s) & \leq \sum_{t=0}^H \frac{\nu v R^{\max}(q + (T - t - 1))^{9/2} \sqrt{\log(5T|\mathcal{Y}|/\delta)}}{\sqrt{NT}} \\ & = \frac{\nu v R^{\max}(T + 1/2T^2)^{9/2} \sqrt{\log(5T|\mathcal{Y}|/\delta)}}{\sqrt{NT}} \end{aligned}$$

□

Proof of Theorem 1. We decompose the regret achieved by the *feasible* estimator as:

$$\tilde{V}_0^{\pi^*} - \tilde{V}_0^{\hat{\pi}^*_{\text{feasible}}} = \tilde{V}_0^{\pi^*} - \tilde{V}_0^{\hat{\pi}^*_{\text{oracle}}} + \tilde{V}_0^{\hat{\pi}^*_{\text{oracle}}} - \tilde{V}_0^{\hat{\pi}^*_{\text{feasible}}}.$$

Theorem 4 established the bound on $\tilde{V}_0^{\pi^*} - \tilde{V}_0^{\hat{\pi}^*_{\text{oracle}}}$. We now show that $\tilde{V}_0^{\hat{\pi}^*_{\text{oracle}}} - \tilde{V}_0^{\hat{\pi}^*_{\text{feasible}}} = o_p(n^{-\frac{1}{2}})$ under the rate assumptions on $\hat{e}, \hat{\mu}$.

Verifying rate double-robustness is standard given arguments in the single-timestep literature. The key step is, as in the proof of Theorem 4, applying the additive error decomposition of Lemma 2.

Then standard single-timestep analysis for doubly-robust policy optimization yields the result [Wager and Athey \(2017\)](#); [Zhou et al. \(2018\)](#). We simply apply our uniform convergence bounds and state the decomposition for completeness.

Let $\hat{P}(Y = y | a)$ denote the estimator for the policy value. Write $\hat{P}^{(1)}(Y = y | a)$ for the estimator evaluated on the first fold using out-of-fold nuisances, and vice-versa.

$$\begin{aligned}\hat{\Gamma}_t^i(y \mid a) &:= \frac{\mathbb{I}[Y_t^i = y] - \hat{\mu}^{-k(i,t)}(y \mid A_t^i, X_t^i)}{\hat{e}_t^{-k(i,t)}(A_t^i \mid X_t^i)} \mathbb{I}[A_t^i = a] + \hat{\mu}^{-k(i,t)}(y \mid a, X_t^i) \\ \hat{\tau}(a, y) &= \frac{1}{n} \sum_{i=1}^N \sum_{a \in \mathcal{A}} \pi(a, X_t^i) \hat{\Gamma}_t^i(y \mid a)\end{aligned}$$

We denote partial sums of the individual contribution to the estimator $\hat{\Gamma}_t^i$, as $\hat{\Gamma}_{\mathcal{I}_1}$, in order to denote the sum over the (sample-split) dataset evaluated with the estimated nuisances, and $\Gamma_{\mathcal{I}_1}$ denote evaluation of the corresponding estimator summands over the oracle nuisance functions.

Due to sample splitting, conditional on a fold $\mathcal{I}_1$, $\hat{\mu}^{(1)}$ can be treated as deterministic. The decomposition used in the proof of Theorem 1, eq. (19) allows us to study regret of the oracle vs. nuisance estimators relative to the true value function $\tilde{V}$. We establish uniformity of the error from using estimated nuisances:

$$\sup_{\pi_{t:T}} \left| \sum_{a \in \mathcal{A}} \sum_{i=1}^n \pi_t(a \mid X_i) (\hat{\Gamma}_t^i(y \mid a) - \Gamma_t^i(y \mid a)) (R(s, a, s') + V_{t+1}^{\pi_{t+1:T}}(s')) \right| = o_p(n^{-\frac{1}{2}}) \quad (21)$$

We decompose the terms as follows, restricting attention to one action, by adding and subtracting $\frac{(\mathbb{I}[Y_t^i=y] - \mu^{-k(i)}(a, X_t^i))}{\hat{e}_t^{-k(i)}(a, X_t^i)}$. To clear up the display we suppress arguments depending on X_i and others where self-evident.

$$\begin{aligned}\mathbb{E}_{N,T}[\pi(a)(\hat{\Gamma}_{\mathcal{I}_1}^i - \Gamma_{\mathcal{I}_1}^i)] &= \\ \mathbb{E}_{N,T} \left[\pi(a) \left(\left(\hat{\mu}^{-k(i,t)} - \mu^{-k(i,t)} \right) + \left(\frac{\mathbb{I}[Y_t = y] - \hat{\mu}^{-k(i,t)}}{\hat{e}_t^{-k(i,t)}} - \frac{\mathbb{I}[Y_t = y] - \mu^{-k(i,t)}}{e_t^{-k(i,t)}} \right) \mathbb{I}[A_t = a] \right) \right] \\ &= \mathbb{E}_{N,T} \left[\pi(a) \left(\hat{\mu}^{-k(i,t)} - \mu^{-k(i,t)} \right) \left(1 - \frac{\mathbb{I}[A_t = a]}{e_t^{-k(i,t)}} \right) (R + V_{t+1}^{\pi_{t+1:T}}(s')) \right] \quad (22)\end{aligned}$$

$$+ \mathbb{E}_{N,T} \left[\pi(a) \left(\mu^{-k(i,t)} - \hat{\mu}^{-k(i,t)} \right) \left(\frac{\mathbb{I}[A_t = a]}{e_t^{-k(i,t)}} - \frac{\mathbb{I}[A_t = a]}{\hat{e}_t^{-k(i,t)}} \right) (R + V_{t+1}^{\pi_{t+1:T}}(s')) \right] \quad (23)$$

$$+ \mathbb{E}_{N,T} \left[\pi(a) (\mathbb{I}[Y_t = y] - \mu^{-k(i,t)}) \left(\frac{1}{\hat{e}_t^{-k(i,t)}} - \frac{1}{e_t^{-k(i,t)}} \right) (R + V_{t+1}^{\pi_{t+1:T}}(s')) \right] \quad (24)$$

Decomposing each of the above terms into foldwise terms, sample splitting implies that conditioning on the other folds implies that μ_1 is a deterministic function.

We apply our Theorem 1 to establish $o_p(n^{-\frac{1}{2}})$ rates on the relevant uniform convergence terms of eqs. (22) and (23).

The term of eq. (22) evaluates to 0 by iterating expectations, using independent errors property from cross-fitting, and well-specification of e which implies $0 = \mathbb{E}_{N,T}[(1 - \mathbb{I}[A_t=a]/e_t^{-k(i,t)}) \mid X]$.

The term of eq. (23) is bounded by Cauchy-Schwarz since $\pi \leq 1$ and by assumption on the sum of rates in Theorem 1, eq. (13)

$$\begin{aligned}\sup_{\pi} (23) &\leq \mathbb{E}_{N,T} \left[\left| \mu^{-k(i,t)} - \hat{\mu}^{-k(i,t)} \right| \left| \frac{\mathbb{I}[A_t = a]}{e_t^{-k(i,t)}} - \frac{\mathbb{I}[A_t = a]}{\hat{e}_t^{-k(i,t)}} \right| \right] \\ &\leq \mathbb{E}_{N,T}[(\mu^{-k(i,t)} - \hat{\mu}^{-k(i,t)})^2]^{\frac{1}{2}} \mathbb{E}_{N,T} \left[\left(\frac{\mathbb{I}[A_t = a]}{e_t^{-k(i,t)}} - \frac{\mathbb{I}[A_t = a]}{\hat{e}_t^{-k(i,t)}} \right)^2 \right]^{\frac{1}{2}} \\ &= o_p(n^{-\frac{1}{2}})\end{aligned}$$

For the term of eq. (24), we apply the bound of Thm. B which surfaces the dependence of the maximal inequality on the behavior of the envelope function, allowing us to leverage consistency assumptions of Theorem 1 to verify that the term is $o_p(n^{-\frac{1}{2}})$.

$$\begin{aligned} \mathbb{E}[\sup_{\pi} (24)] &\leq \mathbb{E}_{N,T}[\pi(a)(\mathbb{I}[Y_t = y] - \mu^{-k(i,t)})(R + V_{t+1}^{\pi_{t+1:T}}(s'))] \mathbb{E}_{N,T} \left[\left| (\hat{e}_t^{-k(i,t)})^{-1} - (e_t^{-k(i,t)})^{-1} \right| \right] \\ &\leq \mathbb{E}_{N,T}[\pi(a)(\mathbb{I}[Y_t = y] - \mu^{-k(i,t)})(R + V_{t+1}^{\pi_{t+1:T}}(s'))] \times \nu \mathbb{E}_{N,T} \left[\left| \hat{e}_t^{-k(i,t)} - e_t^{-k(i,t)} \right| \right] \\ &= o_p(n^{-\frac{1}{2}}) \end{aligned}$$

eq. (21) follows from the above and applying Markov's inequality for the bound for eq. (24). $\square$

Proof of lemma 2. We follow an induction argument. Note that since $\tilde{V}_{T+1} = 0$, the base case, $t = T$, satisfies that

$$\tilde{Q}_T^{\pi_T}(s, \pi_T) - \hat{Q}_T^{\pi_T}(s, \pi_T) = \mathbb{E}[\tilde{R}(s, \pi_T)] - \hat{\mathbb{E}}[\tilde{R}(s, \pi_T)].$$

Now suppose the inductive hypothesis holds for $t + 1$ and consider the case of t .

$$\begin{aligned} \tilde{Q}_t^{\pi}(s, \pi_t) - \hat{Q}_t^{\pi}(s, \pi_t) &= \tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t) + (\tilde{P}\tilde{V}_{t+1}^{\pi})(s, \pi_t) - (\hat{P}\hat{V}_{t+1}^{\pi})(s, \pi_t) \\ &= \tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t) + ((\tilde{P} - \hat{P})\tilde{V}_{t+1}^{\pi})(s, \pi_t) - (\hat{P}(\tilde{V}_{t+1}^{\pi} - \hat{V}_{t+1}^{\pi}))(s, \pi_t) \end{aligned}$$

by a standard performance difference lemma. Since actions in the policy space are themselves policies, the previous analysis for Q functions applies also to V functions,

$$\begin{aligned} \tilde{V}_t^{\pi}(s) - \hat{V}_t^{\pi}(s) &= \mathbb{E}_{\pi} \left[\tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t) + ((\tilde{P} - \hat{P})\tilde{V}_{t+1}^{\pi})(s_t, \pi_t) \mid s_t = s \right] \\ &\quad + \mathbb{E}_{\tilde{M}, \pi_t} \left[\tilde{V}_{t+1}^{\pi}(s_{t+1}) - \hat{V}_{t+1}^{\pi}(s_{t+1}) \mid s_t = s \right] \\ &= \mathbb{E}_{\tilde{M}, \tilde{\pi}} \left[\sum_{t=t'}^T (\tilde{R}(s, \pi_t) - \hat{R}(s, \pi_t)) + (\tilde{P}_t - \hat{P}_t)(\tilde{V}_{t+1}^{\pi})(s_t, \pi_t) \mid s_t = s \right] \end{aligned}$$

where we apply the inductive hypothesis in the first line. $\square$

B.3 Verifying nuisance function rates from pooled episodes

Summary of argument We simply verify that the nuisances can be learned from the pooled episodes at the relevant rates. Note that we learn $\mu(a \mid x)$ from data that can be viewed as the collection of $\{(X_t^i, A_t^i, Y_t^i)\}_{0:T}^n$, and $e(a \mid x)$ from $\{(X_t^i, A_t^i)\}_{0:T}^n$. Throughout the paper, we assume that the behavior policy is stationary (i.e. we do not handle dependent data from history-adapted policies, such as from learning policies, though doing so is a straightforward extension via standard mixing arguments). Hence, the sequence of states is strongly stationary. Nonetheless, dependency across timesteps within an episode is a consequence of S_t, X_t adapted policies (although there is no dependence in the noise of the contextual response across timesteps).

Our argument proceeds as follows. First, we invoke results establishing the mixing properties of the infinite-horizon embedding of episodic finite-horizon MDPs. This general result of Bojun (2020) provides a construction (with a modification to ensure aperiodicity that preserves other properties of the chain) ensures that such an embedding has a steady-state distribution. Therefore, we model nuisance estimation as if we estimate the nuisance functions from a single infinite-horizon trajectory, obtained via the construction of Bojun (2020) by simply sequentially concatenating the episodes (and the perturbations for aperiodicity). In this infinite-horizon embedding, the stationary distribution is given by the (finite-horizon) limiting state-action frequencies, so that the sequence is β -mixing. (Clearly, such an argument generalizes to the case of history-adapted policies with a corresponding dependence on mixing rate of the adapted policy). We invoke results on learning from β -mixing sequences, e.g. Farahmand and Szepesvári (2012), which in particular applies the blocking empirical process argument with a more refined peeling empirical process argument.

B.3.1 Preliminaries

We quote a result on learning from β -mixing data of [Farahmand and Szepesvári \(2012\)](#). The assumptions and result are stated for a generic conditional expectation on a generic process $\{X_t, Y_t\}_{0:T}^{1:N}$. (We of course instantiate the result for the processes $\{X_t^i, A_t^i, Y_t^i\}_{0:T}^{1:N}, \{X_t^i, A_t^i\}_{0:T}^{1:N}$). The estimator of interest is regularized least-squares regression $\hat{m}_n(x) = \text{clip}_{-L,L}(\tilde{m}(x))$ where the final estimates are clipped within a bounded range, and J is a regularization functional.

$$\tilde{m}_n = \operatorname{argmin}_{f \in \mathcal{F}} \left\{ \frac{1}{n} \sum_{i=1}^n |f(X_i) - Y_i|^2 + \lambda_n J^2(f) \right\} \quad (25)$$

Assumption 4 (Exponential Mixing). The process $((X_t, Y_t))_{t=1,2,\dots}$ is an $\mathcal{X} \times \mathbb{R}$ -valued, stationary, exponentially β -mixing stochastic process. In particular, the β -mixing coefficients satisfy $\beta_k \leq \bar{\beta}_0 \exp(-\bar{\beta}_1 k)$, where $\bar{\beta}_0 \geq 0$ and $\bar{\beta}_1 > 0$.

Assumption 5 (Capacity). There exist $C > 0$ and $0 \leq \alpha < 1$ such that for any $u, R > 0$ and all $x_1, \dots, x_n \in \mathcal{X}$

$$\log \mathcal{N}_2(u, \mathcal{B}_R, x_{1:n}) \leq C \left(\frac{R}{u} \right)^{2\alpha}$$

Assumption 6 (Boundedness). There exists $0 < L < \infty$ such that the common distribution of Y_t is such that $|Y_t| \leq L$ almost surely.

Assumption 7 (Realizability). The regression function $m(x) = \mathbb{E}[Y_1 | X_1 = x]$ belongs to the function space $\mathcal{F}$.

Assumption 5 is a mild assumption common in nonparametric statistics; see [\(Györfi et al., 2002\)](#), Lemmas 20.4, 20.6). We verify we may learn the nuisance functions in the following statement, which is a corollary of the construction of [\(Bojun, 2020\)](#), Thm. 2) and the convergence rate of [\(Farahmand and Szepesvári, 2012\)](#), Thm. 5).

Corollary 5 (Nuisance functions). Suppose nuisance function estimator $\hat{e}(a | x)$ is learned from pooled-episode data $\{(X_t^i, A_t^i, Y_t^i)\}_{0:T}^n$, and $\mu(y | a, x)$ from $\{(X_t^i, A_t^i)\}_{0:T}^n$. Let A2-A4 hold, then there exists $c_1, c_2 > 0$ such that for large enough n , for any fixed $\delta \in (0, 1)$, with probability $\geq 1 - \delta$, $\mathbb{E}[(m(X) - \hat{m}(X))^2] \leq c_1 [J^2(m)]^{\frac{\alpha}{1+\alpha}} n^{-\frac{1}{1+\alpha}} \left[\frac{\log(n \vee c_2 / \delta)}{\beta_1} \right]^3$.

Proof of Corollary 2. We first consider the infinite-horizon embedding, $\mathcal{M}^L$, of the finite-horizon MDP. To keep the presentation self-contained, we describe the construction of [Bojun \(2020\)](#). To allow for time-inhomogeneous (but state-stationary) policies, augment the state space with a time state T_t so that the infinite-horizon MDP's state space is (T_t, S_t, X_t) . We embed the finite-horizon MDP by advancing time states in the natural sense, and transitioning from end-of-episodes to initial states,

$$P(t', s_{t+1}, x_{t+1} | t, s_t, x_t, a_t) = \mathbb{I}[t' = t + 1] P(s_{t+1}, x_{t+1} | t, s_t, x_t, a_t) \\ P(kT + 1, s_0, x_{kT+1} | kH, s_{kT}, x_{kT}, a_{kT}) = P(s_0), \forall k \in \mathbb{N}$$

Such a process satisfies the definition of an episodic learning process, e.g. definition 1 of [Bojun \(2020\)](#). Now, to ensure ergodicity due to periodicity of the episodic/finite-horizon structure, [Bojun \(2020\)](#) establishes that a simple ϵ -perturbation recovers aperiodicity by introducing a perturbed MDP, $\mathcal{M}^+$ which simply introduces an auxiliary null state s_{null} . With some ϵ probability, transitions from terminal state (e.g. $T_t \bmod T = 0$) transit to s_{null} before transiting to the initial state and beginning another episode. Theorems 2 and 3 verify ergodicity (evident under aperiodicity) and that the Q functions are identical between $\mathcal{M}^L$ and $\mathcal{M}^+$. Properties of the construction are clear under aperiodicity, see, e.g. [\(Meyn and Tweedie, 2012\)](#).

Clearly, the stationary distribution of the chain is

$$p_{\pi}^{\infty,+}(t, s_t, x_t, a_t) = \frac{1}{T} f_t(x_t) d^{\pi}(s_t) \pi_t(a_t | s_t, x_t),$$

where $d^{\pi}(s_t)$ is the time- t marginal state occupancy distribution under π in the original finite-horizon MDP, $\mathcal{M}$. Convergence to stationarity distribution of $\mathcal{M}^L$ (and hence $\mathcal{M}^+$) is convergence of empirical state-occupancy distributions $d^{\pi}(s_t)$, which converge geometrically. □

B.4 Proofs of structural analysis for dynamic/capacitated pricing

We include the full statement of theorem including the general sufficient condition for $t < T - 2$ that is less interpretable.

Assumption 8 (Reward ordering). $R(1) > R(0)$, without loss of generality.

Assumption 9 (Discrete concavity of value function). V_t is a discrete concave function in s , for all t .

Assumption 10 (Unimodal expected reward). $\tilde{R}(s, \pi)$ is decreasing for suboptimal π (as suboptimal thresholds $\hat{\theta}$ deviate from optimal thresholds).

Assumption 11 (Treatment effect regime.).

$$(2P(Y = 1 \mid 0) - 1) + \mathbb{E}[(2\tau(X) - 1)\mathbb{I}[\Delta\mu(X) > \hat{\theta}_{t+1}(s)]] \leq 0$$

Theorem 3 (full statement for general case, $t < T - 2$). Let $\tau(x) = \eta(1 \mid 1, x) - \eta(1 \mid 0, x)$. For $t = T - 2$, assume assumption 8. Then, for $s > 2$,

$$\mathbb{E}[-\tau(X)\mathbb{I}[\hat{\theta}_{T-1}^* \leq \Delta\mu^*(X) \leq \theta_{T-1}^*]] \geq 0 \implies \hat{\theta}_{T-2}^*(\Delta V_{T-1}^{\hat{\pi}_{T-1}^*, \pi_{T-1}^*}) < \hat{\theta}_{T-2}^*(\Delta V_{T-1}^{\pi_{T-1}^*, \pi_{T-1}^*}).$$

For $t < T - 2$ and $s > 2$, assuming assumptions 8 to 11, if

$$\begin{aligned} & (2V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s-1) - (V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s) + V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s-2)))\mathbb{E}[-\tau(X)\mathbb{I}[\hat{\theta}_{T-1}^* \leq \Delta\mu^*(X) \leq \theta_{T-1}^*]] \\ & + \mathbb{E}_{\pi_t^*(s) - \hat{\pi}_t^*(s)}[R(y) + V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s') \mid s-1] - \mathbb{E}_{\pi_t^*(s-1) - \hat{\pi}_t^*(s-1)}[R(y) + V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s') \mid s-1] \leq 0, \end{aligned}$$

then $\hat{\theta}_t^*(\Delta V_{t+1}^{\hat{\pi}_{t+1}^*, \pi_{t+1}^*}) < \hat{\theta}_t^*(\Delta V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*})$.

These conditions are problem dependent, depending on the true response models. Assumption 9 is satisfied for the specific dynamic pricing example, and Assumption 10 is a common assumption that is satisfied if conditional responses satisfy log-concave noise distributions. Assumption 11 is an assumption about the treatment effect and outcome regime, which would be satisfied in a ‘‘sparse reward, not very large treatment effect’’ regime. We collect some lemmas used for the proof. In this section, we focus on the marginal MDP and for clarity in the notation, omit $\tilde{V}$ and instead write $V_t(S)$ for the value function in the marginal MDP.

Lemma 3 (Discrete concavity of the value function). Suppose $R \geq 0$. For any $t \in [T+1]$ and π , $V_t^\pi(s)$ is (discrete) concave in s .

Lemma 4 (Single-step deviation and difference in value function differences). For $t < T - 1$,

$$\begin{aligned} & \Delta V_t^{\hat{\pi}_t^*, \pi_{t+1}^*}(s) - \Delta V_t^{\pi_t^*, \pi_{t+1}^*}(s) \\ & = (2V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s-1) - (V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s) + V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s-2)))\mathbb{E}[-\tau(X)\mathbb{I}[\hat{\theta}_{T-1}^* \leq \Delta\mu^*(X) \leq \theta_{T-1}^*]] \\ & + \mathbb{E}_{\pi_t^*(s) - \hat{\pi}_t^*(s)}[R(y) + V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s') \mid s-1] - \mathbb{E}_{\pi_t^*(s-1) - \hat{\pi}_t^*(s-1)}[R(y) + V_{t+1}^{\pi_{t+1}^*, \pi_{t+1}^*}(s') \mid s-1] \end{aligned}$$

For $t = T - 1$:

$$\begin{aligned} & \Delta V_{T-1}^{\hat{\pi}_{T-1}^*, \pi_T^*}(s) - \Delta V_{T-1}^{\pi_{T-1}^*, \pi_T^*}(s) \\ & = (2V_T^{\pi_T^*, \pi_T^*}(s-1) - (V_T^{\pi_T^*, \pi_T^*}(s) + V_T^{\pi_T^*, \pi_T^*}(s-2))) \int (\mu(1 \mid 0, x) - \mu(1 \mid 1, x)) \left(\mathbb{I}[\Delta\mu^* > \theta_{T-1}^*] - \mathbb{I}[\Delta\mu^* > \hat{\theta}_{T-1}^*] \right) dFx \end{aligned}$$

Lemma 5 (Value function decomposition). Let $V_t^{\hat{\pi}_t^*, \pi_T^*}(s) = V_t^{\pi_t^*, \pi_T^*}(s) + \delta_t^V(s)$.

$$V_t^{\pi_t^*, \pi_T^*}(s) - V_t^{\hat{\pi}_t^*, \pi_T^*}(s) \tag{26}$$

$$= V_t^{\pi_t^*, \pi_T^*}(s) - V_t^{\hat{\pi}_t^*, \pi_T^*}(s) - \sum_{s', y \mid s} \delta_{t+1}^V(s') \int (\mu(1 \mid 1, x)\mathbb{I}[\Delta\mu > \hat{\theta}_t^*] + \mu(1 \mid 0, x)(1 - \mathbb{I}[\Delta\mu > \hat{\theta}_t^*])) dFx \tag{27}$$

Proof of Theorem 3. First, observe that

$$\hat{\theta}_t^*(\Delta V_{t+1}^{\hat{\pi}_{t+1:T}^*}) \leq \hat{\theta}_t^*(\Delta V_{t+1}^{\pi_{t+1:T}^*}) \iff \Delta V_{t+1}^{\hat{\pi}_{t+1:T}^*} \geq \Delta V_{t+1}^{\pi_{t+1:T}^*} \quad (28)$$

To see this, consider for simplicity $\frac{a+x}{b+x} < \frac{a+y}{b+y}$, where $a = R(0), b = R(1), x = \Delta V^{\hat{\pi}}, y = \Delta V^{\pi^*}$. Note that $a < b$, e.g. $R(0) < R(1)$ and $x > y$ and $\frac{a+x}{b+x} < \frac{a+y}{b+y} \iff (a-b)(y-x) < 0$, i.e. iff $\Delta V_{t+1}^{\hat{\pi}_{t+1:T}^*} > \Delta V_{t+1}^{\pi_{t+1:T}^*}$.

Therefore we will show $\Delta V_t^{\hat{\pi}_{t:T}^*} \geq \Delta V_t^{\pi_{t:T}^*}, \forall t, s > 1$.

We will also establish that for a given timestep, the bias in thresholds is decreasing in the system state (as are the value function differences). To summarize, we consider the following inductive hypotheses:

$$\Delta V_{t+1}^{\hat{\pi}_{t+1:T}^*} \geq \Delta V_{t+1}^{\pi_{t+1:T}^*} \quad (29)$$

$$\theta_{t+1}(s) - \hat{\theta}_{t+1}(s) \leq \theta_{t+1}(s-1) - \hat{\theta}_{t+1}(s-1) \quad (30)$$

$$\Delta V^{\pi_{t+1:T}^*}(s) - \Delta V^{\hat{\pi}_{t+1:T}^*}(s) \leq \Delta V^{\pi_{t+1:T}^*}(s-1) - \Delta V^{\hat{\pi}_{t+1:T}^*}(s-1) \quad (31)$$

We prove the inductive step by assuming the above inductive hypotheses are true for $t' = t+1$ and verifying that this implies they hold for $t' = t$. The main analysis is in verifying $\Delta V_{t+1}^{\hat{\pi}_{t+1:T}^*} \geq \Delta V_{t+1}^{\pi_{t+1:T}^*}$, eq. (29), which requires the other induction hypotheses. We first show the inductive step holds for eqs. (30) and (31) under the induction hypotheses. Note that in the special case of $t = T-2$, only eq. (29) is needed.

Inductive step for eq. (30). To lighten the notation for the following comparisons, we denote $\Delta \hat{V}^*(s) = \Delta V^{\hat{\pi}_{t:T}^*}(s), \Delta V^*(s) = \Delta V^{\pi_{t:T}^*}(s)$.

$$\begin{aligned} \theta_t(s) - \hat{\theta}_t(s) &\leq \theta_t(s-1) - \hat{\theta}_t(s-1) \\ \iff \frac{(R(0) - R(1))(\Delta V^*(s) - \Delta \hat{V}^*(s))}{(R(1) + \Delta V^*(s))(R(1) + \Delta \hat{V}^*(s))} &> \frac{(R(0) - R(1))(\Delta V^*(s-1) - \Delta \hat{V}^*(s-1))}{(R(1) + \Delta V^*(s-1))(R(1) + \Delta \hat{V}^*(s-1))} \\ \iff \frac{\Delta V^*(s) - \Delta \hat{V}^*(s)}{\Delta V^*(s-1) - \Delta \hat{V}^*(s-1)} &< \frac{(R(1) + \Delta V^*(s))(R(1) + \Delta \hat{V}^*(s))}{(R(1) + \Delta V^*(s-1))(R(1) + \Delta \hat{V}^*(s-1))} \end{aligned}$$

where from the second to last line, we use the fact that $R(0) - R(1) < 0$ by assumption (without loss of generality) of Theorem 3 on the ordering of the rewards.

The LHS of the last inequality above is less than 1 by the induction hypothesis on value function differences in states, $\Delta V^*(s) - \Delta \hat{V}^*(s) \leq \Delta V^*(s-1) - \Delta \hat{V}^*(s-1)$.

The RHS is greater than 1 since

$$(R(1) + \Delta V^*(s))(R(1) + \Delta \hat{V}^*(s)) \geq (R(1) + \Delta V^*(s-1))(R(1) + \Delta \hat{V}^*(s-1)),$$

because Lemma 3 implies $\Delta V(s)$ is increasing as s increases, so $\Delta V^*(s) \geq \Delta V^*(s-1)$ and $\Delta \hat{V}^*(s) \geq \Delta \hat{V}^*(s-1)$.

Inductive step for eq. (31).

Equation (31) is equivalent to showing $\Delta V^*(s) - \Delta V^*(s-1) \leq \Delta \hat{V}^*(s) - \Delta \hat{V}^*(s-1)$. First we decompose the difference into the single-stage reward difference term and the policy-induced transitions to next value functions. Having verified eq. (30) for time t , in combination with assumption 10 which implies that greater differences in biased vs. optimal thresholds lead to greater reward suboptimality, implies the single-stage term satisfies the inequality. Verifying the inequality for the difference in value functions term holds by an argument similar to used in showing eq. (30), that the region of integration (even after accounting for differences in thresholds) remains one of positive measure, while the integrand is negative (satisfies the inequality) by the induction hypothesis for eq. (31), for next-time-step value function differences over states.

Inductive step for $\Delta V_{t+1}^{\hat{\pi}_{t+1:T}^*} \geq \Delta V_{t+1}^{\pi_{t+1:T}^*}$:

We first establish the base case by studying some properties of V_T and its differences which simplify the analysis.

$$V_T^{\pi_T^*}(s') - V_T^{\hat{\pi}_T^*}(s') = V_T^{\pi_T^*}(s) - V_T^{\hat{\pi}_T^*}(s), \forall s, s' \geq 1 \quad (32)$$

Therefore $\Delta V_T^{\pi^*}(s) - \Delta V_T^{\hat{\pi}^*}(s) = 0$, for $s > 1$.

Note that θ^* is independent of state. Under Assumption 2, when $s > 2$, $V_T^{\pi^*}(s-1) - V_T^{\hat{\pi}^*}(s-1) = V_T^{\pi^*}(s) - V_T^{\hat{\pi}^*}(s)$, since $V_{T+1}(s) = 0$, since the only non-zero terms are invariant in states. The above follows by rearranging. When $s > 1$, the statement is also true because next-stage value is 0 (independent of downstream estimation error), and it is true by definition for $s = 0$.

To establish the inductive step, we decompose $\Delta V_{t+1}^{\pi_{t+1}^*:T} - \Delta V_{t+1}^{\hat{\pi}_{t+1}^*:T}$. By definition,

$$\Delta V_{t+1}^{\pi_{t+1}^*:T}(s) - \Delta V_{t+1}^{\hat{\pi}_{t+1}^*:T}(s) = (V_{t+1}^{\pi_{t+1}^*:T}(s-1) - V_{t+1}^{\hat{\pi}_{t+1}^*:T}(s-1)) - (V_{t+1}^{\pi_{t+1}^*:T}(s) - V_{t+1}^{\hat{\pi}_{t+1}^*:T}(s)) \quad (33)$$

By lemma 5,

$$\begin{aligned} & \Delta V_{t+1}^{\pi_{t+1}^*:T}(s) - \Delta V_{t+1}^{\hat{\pi}_{t+1}^*:T}(s) \\ &= \Delta V^{\pi_{t+1}^*, \pi_T^*}(s) - \Delta V^{\hat{\pi}_{t+1}^*, \pi_T^*}(s) \\ &+ \sum_{s', y|s} \delta_{t+1}^V(s') \int \left(\mu(1 | 1, x) \mathbb{I}[\Delta\mu > \hat{\theta}_{t+1}^*] + \mu(1 | 0, x)(1 - \mathbb{I}[\Delta\mu > \hat{\theta}_{t+1}^*]) \right) dFx \\ &- \sum_{s', y|s-1} \delta_{t+1}^V(s') \int \left(\mu(1 | 1, x) \mathbb{I}[\Delta\mu > \hat{\theta}_{t+1}^*] + \mu(1 | 0, x)(1 - \mathbb{I}[\Delta\mu > \hat{\theta}_{t+1}^*]) \right) dFx \end{aligned}$$

Let $C(y) = \int \left(\mu(1 | 1, x) \mathbb{I}[\Delta\mu > \hat{\theta}_{t+1}^*] + \mu(1 | 0, x)(1 - \mathbb{I}[\Delta\mu > \hat{\theta}_{t+1}^*]) \right) dFx$.

Then:

$$\Delta V_{t+1}^{\pi_{t+1}^*:T}(s) - \Delta V_{t+1}^{\hat{\pi}_{t+1}^*:T}(s) \quad (34)$$

$$= \underbrace{\Delta V^{\pi_{t+1}^*, \pi_T^*}(s) - \Delta V^{\hat{\pi}_{t+1}^*, \pi_T^*}(s)}_{(1)} + \underbrace{\delta_T^V(s-1)(C(1) - C(0)) + \delta_T^V(s)C(0) - \delta_T^V(s-2)C(1)}_{(2)} \quad (35)$$

We first analyze (1), the value function difference under a single-timestep policy difference. By Lemma 4,

$$\begin{aligned} & \Delta V_t^{\pi_t^*, \pi_{t+1}^*:T}(s) - \Delta V_t^{\hat{\pi}_t^*, \pi_{t+1}^*:T}(s) \\ &= (2V_{t+1}^{\pi_{t+1}^*:T}(s-1) - (V_{t+1}^{\pi_{t+1}^*:T}(s) + V_{t+1}^{\pi_{t+1}^*:T}(s-2)))\mathbb{E}[-\tau(X)\mathbb{I}[\hat{\theta}_t^*(s) \leq \Delta\mu^*(X) \leq \theta_t^*(s)]] \\ &+ \mathbb{E}_{\pi_t^*(s) - \hat{\pi}_t^*(s)}[R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s') | s-1] - \mathbb{E}_{\pi_t^*(s-1) - \hat{\pi}_t^*(s-1)}[R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s') | s-1] \end{aligned}$$

We establish that the last term is equivalent to integrating over an interval, such that the last term is generally positive since the value functions are nonnegative. Therefore, in the general case, negativity of the first term from the sufficient condition when $t = T-2$ isn't completely sufficient. However, the same inequality may hold given that the state-wise differences in biased threshold suboptimality is not too large.

From the inductive hypothesis that $\hat{\theta}_{t+1}(s) < \theta_{t+1}(s), \forall s > 1$, and Lemma 3 which implies that $\Delta V(s)$ is increasing in s (by properties of discrete derivatives of discrete concave functions), we deduce:

$$\hat{\theta}_{t+1}(s) \leq \theta_{t+1}(s) \leq \theta_{t+1}(s-1), \quad \hat{\theta}_{t+1}(s) \leq \hat{\theta}_{t+1}(s-1) \leq \theta_{t+1}(s-1) \quad (36)$$

By the properties in eq. (36) and the induction hypothesis eq. (30), integrating against an increasing function in $\Delta\mu$ is nonnegative:

$$\mathbb{E} \left[g(\Delta\mu) \mathbb{I}[\hat{\theta}_{t+1}^*(s) < \Delta\mu^* < \theta_{t+1}^*(s)] - \mathbb{I}[\hat{\theta}_{t+1}^*(s-1) < \Delta\mu^* < \hat{\theta}_{t+1}^*(s)] \right] \geq 0.$$

We next analyze (2). Note the inductive hypothesis implies $\delta_{t+1}^V(s-1) \leq \delta_{t+1}^V(s)$, e.g. $\delta_{t+1}^V(s)$ is decreasing in s , by rewriting eq. (33) with definition of δ^V . We decompose $C(1) - C(0)$ as:

$$C(1) - C(0) = (2P(Y(0) = 1) - 1) + \int (2(\eta(1 | 1, x) - \eta(1 | 0, x)) - 1) \mathbb{I}[\Delta\mu > \hat{\theta}_{t+1}^*] dFx$$

Therefore, under Assumption [11](#),

$$\begin{aligned} \textcircled{2} &= (\delta_T^V(s-1) - \delta_T^V(s))(C(0) - C(1)) + (\delta_T^V(s-2) - \delta_T^V(s))C(1) \\ &\leq 0 \end{aligned}$$

Simplification when $t = T - 1$. When $t = T - 1$, the sufficient condition is a direct consequence of eq. [\(35\)](#). Since $\delta_T^V(s) = \delta_T^V(s')$, $\forall s, s' \in \mathcal{S}$, $\textcircled{2} = 0$.

Applying Lemma [4](#),

$$\begin{aligned} \Delta V_{T-1}^{\pi_{T-1}^*} (s) - \Delta V_{T-1}^{\hat{\pi}_{T-1}^*} (s) \\ = ((V_T^{\pi_T^*}(s) + V_T^{\pi_T^*}(s-2)) - 2V_T^{\pi_T^*}(s-1))\mathbb{E}[-\tau(X)\mathbb{I}[\hat{\theta}_t^*(s) \leq \Delta\mu^*(X) \leq \theta_t^*(s)]] \end{aligned}$$

The first multiplicative term is negative by Lemma [3](#), concavity of $V^{\pi_T^*}(s)$ in s . $\square$

Proofs of auxiliary lemmas

Proof of Lemma [3](#). This is a structural result of the dynamic pricing problem. We include the proof for completeness but the argument is not novel: we simply verify the adaptation of Theorem 1.18 [\[Gallego et al. \(2019\)\]](#) holds for the contextual setting in this paper.

Note that the difference from that formulation is that randomness is modeled in the transition probabilities, not the arrival rates of consumers, and we express $\Delta V(s) = V(s-1) - V(s)$ as the negative finite difference.

The proof shows concavity of $V_t(s)$ in s by showing that $\Delta V(s)$ is increasing in s ; hence finite differences (discrete derivatives) are decreasing so that the value function V is concave. The argument follows by forward induction on the state space and a sample path argument.

The base case holds by definition of $\Delta V(0) = -\infty$; clearly $\Delta V_t(1) > \Delta V_t(0)$ for any t . The induction hypothesis posits that for some $s+1$, following the optimal policies, $\Delta V_t^{\pi^*}(s')$ is increasing in s' for $s' \leq s$ for all t . We want to show $\Delta V_t^{\pi^*}(s+1) \geq \Delta V_t^{\pi^*}(s)$, or equivalently

$$V_t^{\pi^*}(s+1) + V^{\pi^*}(t, s-1) \leq V^{\pi^*}(s) + V^{\pi^*}(s)$$

We verify the sample path argument of [\[Gallego et al. \(2019\)\]](#) holds in this setting with actions taken in the marginal MDP formulation. We will show:

$$V_t^{\pi^*(s+1)}(s+1) + V_t^{\pi^*(s-1)}(s-1) \leq V^{\pi^*(s+1)}(s) + V^{\pi^*(s-1)}(s)$$

Clearly by suboptimality of the policies optimal at states $s+1, s-1$ for state s , $V^{\pi^*(s+1)}(s) + V^{\pi^*(s-1)}(s) \leq V^{\pi^*(s)}(s) + V^{\pi^*(s)}(s)$ so showing the above inequality is sufficient to verify the inductive step.

The sample path argument tracks the usage of the *suboptimal* policies of the left-hand-side original-state $s+1$ and $s-1$ systems, $\pi^*(s+1), \pi^*(s-1)$ for the right-hand-side state- s system, until one of the following cases: $t = T$ (time runs out); at some $t' \geq t$ the difference in inventories of the state- $s+1$ and state $s-1$ systems drops to 1, or the state of the original state $s-1$ system drops to 0. Then the optimal policies for the system state are followed thereafter.

Case 1: Use $\pi^*(s+1), \pi^*(s-1)$ for the two state s systems, respectively, until the end of selling horizon.

The realized revenues by following the same randomness sample path and same action policies are identical by following the same policies.

Case 2: At some $t' \geq t$ the difference in inventories of the state $s+1$ and state $s-1$ systems drops to 1.

Up to this stopping time, the realized revenues are identical by the stopping path argument. Because the transition realizations are identical, then the right-hand-side systems have the same state space as the left-hand-side systems, since under $\pi^*(s+1)$, $s-s'$ items sold while under $\pi^*(s-1)$, $s-s'-1$ items had sold. At some $t' \geq t$, at some state $s' \leq s$, following optimal policies thereafter, the LHS value functions are given by

$$V_{t'}^{\pi^*(s'+1)}(s'+1) + V_{t'}^{\pi_{t+1:T}^*}(s'+1) \geq V^{\pi^*(s')}(s') + V_{t'}^{\pi_{t+1:T}^*}(s') + V_{t'}^{\pi^*(s'+1)}(s'+1)$$

so that the remaining optimal expected revenues are identical.

Case 3: For some $t' < T$, original state $s - 1$ system stocks out, e.g. has $s_{t'} = 0$.

At this point, following policy $\pi_t^*(s+1)$ yields a state s' such that $1 < s' \leq x+1$ (the first inequality holds because otherwise we would be in case 2), so that the states at this timepoint are $(t', s'), (t', 0)$. By the sample path argument, identical amounts of goods have sold so the states of the right-hand-side systems are $(t', s' - 1), (t', 1)$. By the inductive hypothesis, $\Delta V_{t'}^{\pi_{t'}^*(s')}(s') \geq \Delta V_{t'}^{\pi_{t'}^*(1)}(1)$ for all $s' \leq s+1$ and all $t' \leq t$. We verify the downstream revenues are at least as high for the right hand systems:

$$V_{t'}^{\pi^*}(s') + V_{t'}(0) \leq V_{t'}(s' - 1) + V_{t'}(1)$$

The inductive hypothesis verifies that $\Delta V(s') \geq \Delta V(1)$, which verifies the above. $\square$

Proof of Lemma 4. Since we restrict attention to single-timestep deviations (following the optimal policy after), we can write $\hat{\pi}^*(\Delta V_{t+1}^{\pi_{t+1}^*:T}(s))$ in terms of a single threshold based on the optimal value function difference, $\theta^* = \theta^*(\Delta V_{t+1}^{\pi_{t+1}^*:T}(s))$.

$$\begin{aligned} V_t^{\pi_t^*, \pi_{t+1}^*:T}(s) - V_t^{\hat{\pi}_t^*, \pi_{t+1}^*:T}(s) &= \mathbb{E}_{\pi_t^*(s) - \pi_t^*(s)}[R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s') \mid s] \\ &= \sum_{s', y \mid s} \int \mu(y \mid 1, x)(R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s')) \left(\mathbb{I}[\Delta\mu^* > \theta_t^*(s)] - \mathbb{I}[\Delta\mu^* > \hat{\theta}_t^*(s)] \right) dFx \\ &\quad + \int \mu(y \mid 0, x)(R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s')) \left(\mathbb{I}[\Delta\mu^* < \theta_t^*(s)] - \mathbb{I}[\Delta\mu^* < \hat{\theta}_t^*(s)] \right) dFx \\ &= \sum_{s', y \mid s} \int \left(\mu(y \mid 1, x)(R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s')) - \mu(y \mid 0, x)(R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s')) \right) \\ &\quad \left(\mathbb{I}[\Delta\mu^* > \theta_t^*(s)] - \mathbb{I}[\Delta\mu^* > \hat{\theta}_t^*(s)] \right) dFx \end{aligned}$$

The claim follows from the above simplification and by expanding the definition,

$$\Delta V_{T-1}^{\hat{\pi}_{T-1}^*, \pi_T^*}(s) - \Delta V_{T-1}^{\pi_{T-1}^*, \pi_T^*}(s) = (V_t^{\pi_t^*, \pi_{t+1}^*:T}(s) - V_t^{\hat{\pi}_t^*, \pi_{t+1}^*:T}(s)) - (V_t^{\pi_t^*, \pi_{t+1}^*:T}(s-1) - V_t^{\hat{\pi}_t^*, \pi_{t+1}^*:T}(s-1)),$$

and algebraic manipulation of the resulting expression. The statement for $t = T - 1$ follows since θ_{T-1} is the same for all states, so that the reward differences also cancel out when the differences of ΔV are considered.

For $t < T - 1$, relative to the simplification of $T - 1$ we obtain an additional term that arises from the differences in $\theta_{t+1}(s) - \hat{\theta}_{t+1}(s)$ for different states s :

$$\mathbb{E}_{\pi_t^*(s) - \pi_t^*(s)}[R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s') \mid s] - \mathbb{E}_{\pi_t^*(s) - \pi_t^*(s)}[R(y) + V_{t+1}^{\pi_{t+1}^*:T}(s') \mid s-1]$$

$\square$

Proof of lemma 5.

$$\begin{aligned} V_t^{\pi_t^*:T}(s) - V_t^{\hat{\pi}_t^*:T}(s) &= \sum_{s', y \mid s} \int (\mu(1 \mid 1, x)(R + V_{t+1}^{\pi_{t+1}^*:T}(s')) \mathbb{I}[\Delta\mu > \theta_t^*] + \mu(1 \mid 0, x)(R + V_{t+1}^{\pi_{t+1}^*:T}(s')))(1 - \mathbb{I}[\Delta\mu > \theta_t^*]) dFx \\ &\quad - \sum_{s', y \mid s} \int (\mu(1 \mid 1, x)(R + V_{t+1}^{\pi_{t+1}^*:T}(s') + \delta_{t+1}^V(s')) \mathbb{I}[\Delta\mu > \hat{\theta}_t^*] + \mu(1 \mid 0, x)(R + V_{t+1}^{\pi_{t+1}^*:T}(s') + \delta_{t+1}^V(s')) \\ &\quad \cdot (1 - \mathbb{I}[\Delta\mu > \hat{\theta}_t^*]) dFx \end{aligned}$$

and collect terms corresponding to $V_t^{\pi_t^*:T}(s) - V_t^{\hat{\pi}_t^*:T}(s)$. $\square$

C Discussion

C.1 Further Related Work

Algorithmic analysis under known distributions. Algorithmic analysis, building on online/approximation algorithms and approximate dynamic programming also requires known demand distributions, hence is complementary [Gallego et al. \(2019\)](#). Our approach is particularly beneficial in handling high-dimensional context variables X_t . Naive extensions of these approaches, for example applying them to an MDP with state aggregation on X_t , incurs statistical bias in general due to discretization. On the other hand, using *action-history dependent* policies achieves stronger regret guarantees in recent work, e.g., that include resolving (model-predictive control) [Bumpensanti and Wang \(2020\)](#). In contrast, we restrict to state- and time-dependent, but history-independent, policy specifications.

Off-policy policy learning leveraging off-policy evaluation. We also compare to backwards-recursive off-policy learning approaches in the dynamic treatment regime literature. In some sense, the DTR/longitudinal causal inference is the opposite of our setting: the difficulty arises from *longitudinal dynamics of the same individual*. [Zhang et al. \(2013\)](#) studied an AIPW estimator in the dynamic treatment regime setting but handle policy-dependent nuisances by approximating with a Q function optimized by another method. [Zhao et al. \(2015\)](#) proposes “backwards outcome-weighted learning” which considers backwards induction on an inverse-propensity weighted estimator that conducts importance sampling in the space of trajectories. Their direct consistency analysis of the backwards induction incurs exponential dependence on horizon.

Clarification to other settings. [Emek et al. \(2020\)](#) introduces “stateful online learning”, a version of online adversarial learning with state information, but their setting is different. In particular, they focus on MDPs with deterministic transitions and assume bounded-loss simulatability from any state, focusing on the adversarial setting. Our focus on *offline contextual decision-making* with state information is different from the contextual MDP model where contexts index MDP models themselves.

C.2 Additional examples of stateful problems

Example 5 (Multi-item network revenue management). Multi-item network revenue management is easily modeled as a modification of Example [1](#) with additional outcomes (products). Consider a setting with J different products and K many resources, so that $M \in \mathbb{R}^{K \times J}$ is the resource consumption matrix, where M_{ij} describes how much of resource i product j requires. Denote the event $\mathbb{I}[s \text{ feas. for } j] := \prod_{i \in [K]} \mathbb{I}[M_{ij} < s_i]$, which describes the event that the state variable s is feasible to produce product j .

We suppose a joint distribution on (X, Z) , e.g. we have exogenous context arrivals and exogenous $Z \mid X$ product types (which may be conditional on the context in the most general case). Therefore at each timestep we sell at most one product at a time.

The multiple product $Q_t(s, x, j, a)$ function on the expanded state space (including product arrival type) is analogous. In this case, the context-marginalized value notation, $\bar{V}_t^\pi(S)$, is overloaded: it now marginalizes over the joint distribution of contexts and product types.

Example 6 (Pricing and repositioning). We adapt a simplified example of setting rental price for vehicles at beginning of each period in a finite (or possibly infinite) planning horizon to a contextual setting [El Shar and Jiang \(2020\)](#). Repositioning is achieved by setting prices to induce directional demand. Discrete state space $\mathcal{S}$ denotes the number of cars at a station, with $\bar{s}$ the maximum number of cars in vehicle sharing system. Between locations there is a known origin-destination transition probability ϕ_{ij} . $Y_{ik,t+1}$ is a random variable taking values in $[N]$ that represents the random destination of customer k at station i ; observed at the beginning of period $t + 1$. Uncontextually, $Y_{ik,t+1} = j$ w.p. ϕ_{ij} . Contextually, we consider (X_t, O_t, D_t) exogeneous covariate and origin-destination request, and the individual demand is a binary outcome in response to price, $Y(p_{it})$. To determine the cost function, let $\ell(i, j)$ be the distance from station i to j , and consider a lost sales unit cost $\rho_i, i \in [N]$. The decision vector $\mathbf{p}_t = \{p_{it} \in [\underline{p}_i, \bar{p}_i], \forall i \in [N]\}$ sets prices for each station. To instantiate the key assumption in this setting, our stateful formulation holds if we believe that the underlying system state S_t is not a confounder because it does not affect whether or not an *individual* demand arrival responds to price.

$$\begin{aligned}
 V^*(S_t, X_t, (O_t, D_t)) = \\
 \max_{\mathbf{p}_t} \underbrace{\mathbb{E}[p_{it}\mathbb{I}[Y = 1] \mid X_t, p_{it}, O_t, D_t]\ell(O_t, D_t)}_{\text{spatial pricing}} - \underbrace{\rho_{O_t}\mathbb{I}[S_t(O_t) = 0]}_{\text{lost sales penalty}} + \gamma\mathbb{E}[\tilde{V}^*(S_{t+1}) \mid p_{it}, X_t, O_t, D_t]
 \end{aligned}$$