



Delft University of Technology

Document Version

Final published version

Licence

Dutch Copyright Act (Article 25fa)

Citation (APA)

Kolarijani, M. A. S., & Esfahani, P. M. (2026). From Optimization to Control: Quasi-Policy Iteration. *IEEE Transactions on Automatic Control*, 71(6), 3880-3893. <https://doi.org/10.1109/TAC.2026.3652902>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

In case the licence states "Dutch Copyright Act (Article 25fa)", this publication was made available Green Open Access via the TU Delft Institutional Repository pursuant to Dutch Copyright Act (Article 25fa, the Taverne amendment). This provision does not affect copyright ownership.
Unless copyright is transferred by contract or statute, it remains with the copyright holder.

Sharing and reuse

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

This work is downloaded from Delft University of Technology.

From Optimization to Control: Quasi-Policy Iteration

Mohamad Amin Sharifi Kolarijani  and Peyman Mohajerin Esfahani 

Abstract—Recent control algorithms for Markov decision processes (MDPs) have been designed using an implicit analogy with well-established optimization algorithms. In this article, we adopt the quasi-Newton method (QNM) from convex optimization to introduce a novel control algorithm coined as quasi-policy iteration (QPI). In particular, QPI is based on a novel approximation of the “Hessian” matrix in the policy iteration algorithm, which exploits two linear structural constraints specific to MDPs and allows for the incorporation of prior information on the transition probability kernel. While the proposed algorithm has the same computational complexity as value iteration, it exhibits an empirical convergence behavior similar to that of QNM with a low sensitivity to the discount factor.

Index Terms—Dynamic programming, Markov decision processes (MDPs), optimization algorithms, quasi-Newton methods (QNM), reinforcement learning.

I. INTRODUCTION

THE problem of control, or the decision-making problem as it is also known within the operations research community, has been the subject of much research since the introduction of the Bellman principle of optimality in the late 1950s [1]. In particular, the connection between control algorithms for Markov decision processes (MDPs) and optimization algorithms has been noticed since the late 1970s [2]: value iteration (VI) [1] can be seen as an instance of gradient descent (GD) algorithm, and policy iteration (PI) [3] is an instance of the Newton method (NM).

More recent works have used the relationship mentioned above to develop new control algorithms, with faster convergence and/or lower complexity, inspired by their counterparts for solving optimization problems [4], [5]. In the case of *model-based* (a.k.a. planning, with access to the model of the MDP)

algorithms, the combination of the VI algorithm with Polyak momentum [6], Nesterov acceleration [7], and Anderson acceleration [8] has been explored in [9] and [10]. More recently, Halpern anchoring acceleration [11] has been used to introduce the Anchored VI algorithm [12], which in particular exhibits an $\mathcal{O}(1/k)$ -rate for large values of discount factor and even for $\gamma = 1$. Also of notice is the Generalized Second-Order VI algorithm [13], which applies NM on a *smoothed* version of the Bellman operator. For *model-free* (a.k.a. learning, with access to samples from the MDP) algorithms, Speedy Q-Learning (SQL) [14], Momentum Q-learning (QL) [15], and Nesterov Stochastic Approximation [16] are among the algorithms that use the idea of momentum for accelerating stochastic GD algorithm [17], [18], [19], [20] in order to achieve a better rate of convergence compared to standard QL. Moreover, the Zap QL (ZQL) algorithm [21] can be thought of as a second-order learning algorithm which was inspired by the stochastic Newton–Raphson algorithm [22].

We note that all of the aforementioned connections have been focused on *finite* state-action MDPs. When it comes to infinite (continuous) state-action spaces, except in special cases, such as linear–quadratic regulators, one needs to resort to finite-dimensional approximation techniques for computational purposes. This approximation may be at the modeling level by aggregation (discretization) of the state and action spaces, which readily falls into the finite state-action MDP setting [23], [24]. Alternatively, one may directly approximate the value function via finite parametrization and minimizing (a proxy of) the residual of its fixed-point characterization based on the Bellman principle of optimality [25], [26]. Examples of such include linear parameterization [27], [28], or nonlinear parameterization with, for instance, neural network architectures [29], [30], [31] or max-plus approximation [32], [33], [34], [35], [36], [37]. We also note that there is an alternative characterization of the original function as the solution to an infinite-dimensional linear program [38], paving the way for approximation techniques via finite tractable convex optimization [39], [40], [41]. With this view of the literature, it is worth noting that one can cast almost all of these approximation techniques as the solution to a finite-dimensional fixed-point or convex optimization problem.

Main contribution: Motivated by the connection between control and optimization algorithms, we adopt the quasi-Newton method (QNM) from convex optimization to introduce the quasi-policy iteration (QPI) algorithm with the following distinct features.

Received 26 August 2025; accepted 24 December 2025. Date of publication 12 January 2026; date of current version 1 June 2026. This work was supported by the European Research Council (ERC) under Grant TRUST-949796 and in part by the NSERC Discovery under Grant RGPIN-2025-06544. Recommended by Associate Editor A. A. Malikopoulos. (Corresponding author: Mohamad Amin Sharifi Kolarijani.)

Mohamad Amin Sharifi Kolarijani is with the Delft Center for Systems and Control, Delft University of Technology, 2628, CD Delft, The Netherlands (e-mail: M.A.SharifiKolarijani@tudelft.nl).

Peyman Mohajerin Esfahani is with the Delft Center for Systems and Control, Delft University of Technology, 2628, CD Delft, The Netherlands, and also with the Department of Mechanical & Industrial Engineering, University of Toronto, Toronto M5S 3G8, ON, Canada (e-mail: P.MohajerinEsfahani@utoronto.ca).

Digital Object Identifier 10.1109/TAC.2026.3652902

- 1) *Hessian approximation via structural information*: QPI is based on a novel approximation of the “Hessian” matrix in the PI algorithm by exploiting two linear structural constraints specific to MDPs and by allowing for the incorporation of prior information on the transition probability kernel of the MDP (Theorem 3.1). In the special case of incorporating a uniform prior for the transition kernel, QPI can be viewed as a modification of the standard VI using two novel directions with adaptive step-sizes (Corollary 3.2).
- 2) *Convergence and sensitivity to discount factor*: The per-iteration computational complexity of QPI is the same as VI, and its linear convergence can be guaranteed by safeguarding against standard VI (Theorem 3.1) or backtracking (Lemma 3.3). However, in our numerical simulations with random and structured MDPs, QPI exhibits an empirical behavior similar to QNMs with the convergence rate being less sensitive to the discount factor (see Fig. 1).
- 3) *Local superlinear convergence*: We provide a modified implementation of QPI that also incorporates the secant-type constraints in approximation of the Hessian to guarantee the local superlinear convergence (Theorem 3.4).
- 4) *Extension to model-free control (a.k.a. RL)*: We also introduce the quasi-policy learning (QPL) algorithm, the stochastic version of QPI, as a novel model-free algorithm with guaranteed convergence and the same per-iteration complexity as standard QL algorithm (Theorem 3.5).

The rest of this article is organized as follows. In Section II, we describe the optimal control problem of MDPs. In Section III, we introduce and analyze the model-based QPI algorithm and its model-free extension, the QPL algorithm. All the technical proofs are provided in Section IV. The performance of these algorithms is then compared with multiple control algorithms via extensive numerical experiments in Section V. Finally, Section VI concludes this article by providing some final remarks.

Notations: For a vector $v \in \mathbb{R}^n$, we use $v(i)$ and $[v](i)$ to denote its i th element. Similarly, $M(i, j)$ and $[M](i, j)$ denote the element in row i and column j of the matrix $M \in \mathbb{R}^{m \times n}$. We use $\cdot^\top$ to denote the transpose of a vector/matrix. We use $\|\cdot\|_2$ and $\|\cdot\|_\infty$ to denote the 2-norm and ∞ -norm of a vector, respectively. We use $\|\cdot\|_2$ and $\|\cdot\|_F$ for the induced 2-norm and the Frobenius norm of a matrix, respectively. Let $x \sim \mathbb{P}$ be a random variable with distribution $\mathbb{P}$. We particularly use $\hat{x} \sim \mathbb{P}$ to denote a sample of the random variable x drawn from the distribution $\mathbb{P}$. We use $\mathbf{1}$ and $\mathbf{0}$ to denote the all-one and all-zero vectors, respectively. I and $E = \mathbf{1}\mathbf{1}^\top$ denote the identity and all-one matrices, respectively. We denote the i th unit vector by e_i , that is, the vector with its i th element equal to 1 and all other elements equal to 0.

II. OPTIMAL CONTROL OF MDPs

A common formulation of the control problem relies on the concept of MDPs. MDPs are a powerful modeling framework for stochastic environments that can be controlled to minimize some measure of cost. An MDP is a tuple $(\mathcal{S}, \mathcal{A}, \mathbb{P}, c, \gamma)$, where $\mathcal{S}$ and $\mathcal{A}$ are the state space and action space, respectively.

The transition kernel $\mathbb{P}$ encapsulates the state dynamics: for each triplet $(s, a, s^+) \in \mathcal{S} \times \mathcal{A} \times \mathcal{S}$, it gives the probability $\mathbb{P}(s^+|s, a)$ of the transition to state s^+ given that the system is in state s and the chosen control is a . The cost function $c : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$, bounded from below, represents the cost $c(s, a)$ of taking the control action a while the system is in state s . The discount factor $\gamma \in (0, 1)$ can be seen as a tradeoff parameter between short- and long-term costs. *In this study, we consider tabular MDPs with a finite state-action space. In particular, we take $\mathcal{S} = \{1, 2, \dots, n\}$ and $\mathcal{A} = \{1, 2, \dots, m\}$. This, in turn, allows us to treat functions $f : \mathcal{S} \rightarrow \mathbb{R}$ and $g : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ as vectors $f \in \mathbb{R}^{|\mathcal{S}|} = \mathbb{R}^n$ and $g \in \mathbb{R}^{|\mathcal{S} \times \mathcal{A}|} = \mathbb{R}^{nm}$ —in the latter case, we are considering a proper 1-to-1 mapping $\mathcal{S} \times \mathcal{A} \rightarrow \{1, 2, \dots, nm\}$.*

Let us now fix a control policy $\pi : \mathcal{S} \rightarrow \mathcal{A}$, i.e., a mapping from states to actions. The stage cost of the policy π is denoted by $c^\pi \in \mathbb{R}^{|\mathcal{S}|} = \mathbb{R}^n$ with elements $c^\pi(s) = c(s, \pi(s))$ for $s \in \mathcal{S}$. The transition (probability) kernel of the resulting Markov chain under the policy π is denoted by $\mathbb{P}^\pi$, where $\mathbb{P}^\pi(s^+|s) = \mathbb{P}(s^+|s, \pi(s))$ for $s, s^+ \in \mathcal{S}$. We also define the matrix $P^\pi \in \mathbb{R}^{|\mathcal{S}| \times |\mathcal{S}|} = \mathbb{R}^{n \times n}$, with elements $P^\pi(s, s^+) := \mathbb{P}^\pi(s^+|s)$ for $s, s^+ \in \mathcal{S}$, to be the corresponding transition (probability) matrix. The value of a policy is the expected, discounted, accumulative cost of following this policy over an infinite-horizon trajectory: For the policy π , we define the value function $v^\pi \in \mathbb{R}^{|\mathcal{S}|} = \mathbb{R}^n$ with elements

$$v^\pi(s) := \mathbb{E}_{s_{t+1} \sim \mathbb{P}^\pi(\cdot|s_t)} \left[\sum_{t=0}^{\infty} \gamma^t c^\pi(s_t) \mid s_0 = s \right]$$

and the Q-function $q^\pi \in \mathbb{R}^{|\mathcal{S} \times \mathcal{A}|} = \mathbb{R}^{nm}$ with elements

$$q^\pi(s, a) := c(s, a) + \gamma \mathbb{E}_{s^+ \sim \mathbb{P}(\cdot|s, a)} [v^\pi(s^+)]$$

so that we also have $v^\pi(s) = q^\pi(s, \pi(s))$ for each $s \in \mathcal{S}$. Given a value function v , let us also define $\pi_v : \mathcal{S} \rightarrow \mathcal{A}$ by

$$\pi_v(s) \in \arg\min_{a \in \mathcal{A}} \{c(s, a) + \gamma \mathbb{E}_{s^+ \sim \mathbb{P}(\cdot|s, a)} [v(s^+)]\}$$

to be the *greedy policy w.r.t. v*. Similarly, for a Q-function q , define $\pi_q : \mathcal{S} \rightarrow \mathcal{A}$ by

$$\pi_q(s) \in \arg\min_{a \in \mathcal{A}} q(s, a)$$

to be the *greedy policy w.r.t. q*. The problem of interest is to control the MDP optimally, that is, to find the optimal policy π^* with the optimal value/Q-function

$$v^* = \min_{\pi} v^\pi, \quad q^* = \min_{\pi} q^\pi \quad (1)$$

so that the expected, discounted, infinite-horizon cost is minimized. Let us also note that the optimal policy, i.e., the minimizer of the preceding optimization problems, is the greedy policy w.r.t. v^* and q^* , that is, $\pi^* = \pi_{v^*} = \pi_{q^*}$.

Interestingly, the optimal value/Q-function introduced in (1) can be equivalently characterized as the fixed-point of the corresponding Bellman operators. This fixed-point characterization is the basis for the class of VI algorithms. To be precise, we have $v^* = T(v^*)$, i.e.,

$$v^*(s) = [T(v^*)](s) \quad \forall s \in \mathcal{S} \quad (2)$$

where $T : \mathbb{R}^{|\mathcal{S}|} \rightarrow \mathbb{R}^{|\mathcal{S}|}$ is the *Bellman operator* given by

$$[T(v)](s) := \min_{a \in \mathcal{A}} \{c(s, a) + \gamma \mathbb{E}_{s^+ \sim \mathbb{P}(\cdot|s, a)} [v(s^+)]\}. \quad (3)$$

Similarly, we have $q^* = \mathbb{E}[\hat{T}(q^*)]$, i.e.,

$$q^*(s, a) = \mathbb{E} \left[[\hat{T}(q^*)](s, a) \right] \quad \forall (s, a) \in \mathcal{S} \times \mathcal{A} \quad (4)$$

where $\hat{T} : \mathbb{R}^{|\mathcal{S} \times \mathcal{A}|} \rightarrow \mathbb{R}^{|\mathcal{S} \times \mathcal{A}|}$ is the *sampled* Bellman operator given by¹

$$[\hat{T}(q)](s, a) := c(s, a) + \gamma \min_{a^+ \in \mathcal{A}} q(\hat{s}^+, a^+) \quad (5)$$

with $\hat{s}^+ \sim \mathbb{P}(\cdot | s, a)$ being a *sample* of the next state drawn from the distribution $\mathbb{P}(\cdot | s, a)$ for the pair (s, a) . We note that the expectation in (4) is w.r.t. $\hat{s}^+$.

III. QUASI-POLICY ITERATION

In this section, we develop and analyze the QPI algorithm for model-based control problems and its extension, the QPL algorithm, for model-free control problems.² These algorithms, as the name suggests, are inspired by QNMs from convex optimization. For this reason, we begin with a brief overview of the connection between optimization and control algorithms, as well as QNMs, to motivate the algorithms developed in this study.

We note that for tabular MDPs with $\mathcal{S} = \{1, \dots, n\}$ and $\mathcal{A} = \{1, \dots, m\}$, we have $v \in \mathbb{R}^n$ and $q \in \mathbb{R}^{nm}$ for the value function and the Q-function, respectively. Correspondingly, we have $T : \mathbb{R}^n \rightarrow \mathbb{R}^n$ with

$$T(v) = \sum_{s \in \mathcal{S}} [T(v)](s) \cdot e_s$$

where $e_s \in \mathbb{R}^n$ is the unit vector for the state $s \in \mathcal{S}$, and also $\hat{T} : \mathbb{R}^{nm} \rightarrow \mathbb{R}^{nm}$ with

$$\hat{T}(q) = \sum_{(s,a) \in \mathcal{S} \times \mathcal{A}} [\hat{T}(q)](s, a) \cdot e_{(s,a)}$$

where $e_{(s,a)} \in \mathbb{R}^{nm}$ is the unit vector for the state-action pair $(s, a) \in \mathcal{S} \times \mathcal{A}$.

A. Optimization Versus Control

Two classical algorithms for solving the optimal control problem of MDPs are VI and PI, which are instances of GD and NM, respectively, as we mentioned above. To be precise, consider the unconstrained minimization problem

$$\min_{x \in \mathbb{R}^\ell} f(x)$$

where $f : \mathbb{R}^\ell \rightarrow \mathbb{R}$ is twice continuously differentiable and strongly convex, and let $g(x) = \nabla f(x)$ and $H(x) = \nabla^2 f(x)$ be the gradient and Hessian of f evaluated at x , respectively. Then, by defining $g(v) := v - T(v)$, one can see that the VI update rule [1]

$$v_{k+1} = T(v_k) = v_k - g(v_k)$$

¹Strictly speaking, the provided sampled Bellman operator is the empirical version of the Bellman operator for the Q-function, given by $[\hat{T}(q)](s, a) := c(s, a) + \gamma \mathbb{E}_{s^+ \sim \mathbb{P}(\cdot | s, a)} [\min_{a^+ \in \mathcal{A}} q(s^+, a^+)]$ for $(s, a) \in \mathcal{S} \times \mathcal{A}$.

²In this article, the terminologies of “model-free” and “model-based” indicate the *available information (oracle)*, i.e., whether we have access to the model or only the system trajectory (samples). We note that this is different from the common terminologies in the RL literature where these terms refer to the *solution approach*, i.e., whether we identify the model along the way (model-based RL) or directly solve the Bellman equation to find the value function (model-free RL).

resembles the GD update rule

$$x_{k+1} = x_k - \alpha_k g(x_k)$$

where α_k is the step-size. Similarly, by defining $H(v) := I - \gamma P^{\pi_v}$, the PI update rule [42, Prop. 6.5.1]

$$v_{k+1} = (I - \gamma P^{\pi_{v_k}})^{-1} c^{\pi_{v_k}} = v_k - [H(v_k)]^{-1} g(v_k)$$

resembles the NM update rule

$$x_{k+1} = x_k - \alpha_k [H(x_k)]^{-1} g(x_k)$$

where α_k is again the step-size.

The VI algorithm converges linearly with rate γ owing to the fact that the operator T is a γ -contraction in the ∞ -norm [42, Prop. 6.2.4]. The PI algorithm, on the other hand, outputs the optimal policy in a finite number of iterations [42, Thm. 6.4.2]. Moreover, the algorithm has a local *quadratic* rate of convergence when initiated in a small enough neighborhood around the optimal solution [43], [44]. The faster convergence of PI compared to VI, however, comes with a higher per-iteration computational complexity: The per-iteration complexities of VI and PI are $\mathcal{O}(n^2 m)$ and $\mathcal{O}(n^2 m + n^3)$, respectively. The extra $\mathcal{O}(n^3)$ complexity is due to the policy evaluation step, i.e., solving a linear system of equations; see also the matrix inversion in the characterization above. Not surprisingly, we see the same convergence-complexity tradeoff between GD and NM. While NM has a better convergence rate compared to GD, it suffers from a higher per-iteration computational cost. GD, with a proper choice of step-size, converges *linearly* [45, Thm. 3.12] with $\mathcal{O}(\ell)$ per-iteration complexity (disregarding the complexity of gradient oracle). On the other hand, NM, with a proper choice of step-size, has a local *quadratic* convergence rate [45, Thm. 5.3] with $\mathcal{O}(\ell^3)$ per-iteration complexity, assuming direct inversion (and disregarding the complexity of gradient and Hessian oracles).

QNMs are a class of methods that allow for a tradeoff between computational complexity and (local) convergence rate. To do so, these methods use a Newton-type update rule

$$x_{k+1} = x_k - \alpha_k \tilde{H}_k^{-1} \nabla f(x_k)$$

where $\tilde{H}_k$ is an *approximation* of the true Hessian $\nabla^2 f(x_k)$ at iteration k . Different QNMs use different approximations of the Hessian. A generic approximation scheme in QNMs is

$$\begin{aligned} \tilde{H}_k &= \arg \min_{H \in \mathbb{R}^{\ell \times \ell}} \|H - H_{\text{prior}}\|_F^2 \\ \text{s.t. } H r_i &= b_i, \quad i = 1, \dots, j \end{aligned} \quad (6)$$

which minimizes the distance (in Frobenius norm) to a given prior H_{prior} subject to j (≥ 1) linear constraints specified by $r_i, b_i \in \mathbb{R}^\ell$. This leads to the approximation $\tilde{H}_k$ being a rank- j update of the prior H_{prior} , i.e.,

$$\tilde{H}_k = H_{\text{prior}} + (B - H_{\text{prior}} R)(R^\top R)^{-1} R^\top \quad (7)$$

where $R = (r_1, \dots, r_j)$, $B = (b_1, \dots, b_j) \in \mathbb{R}^{\ell \times j}$. Hence, $\tilde{H}_k^{-1}$ can be easily computed based on H_{prior}^{-1} using the Woodbury formula. Different choices of the prior and the linear constraints in the generic approximation scheme above lead to different QNMs. For example, by choosing the so-called *secant conditions* with $r_i = x_{k-i+1} - x_{k-i}$ and $b_i = \nabla f(x_{k-i+1}) - \nabla f(x_{k-i})$ as linear constraints and $H_{\text{prior}} = I$ as the prior, we

derive Anderson mixing with memory j [8], while by using a single secant condition with $j = 1$ and choosing $H_{\text{prior}} = \tilde{H}_{k-1}$ as the prior, we derive QNM with Broyden approximation [46].

In what follows, we use a similar idea and propose the QPI algorithm by incorporating a computationally efficient approximation of the ‘‘Hessian’’ $H = I - \gamma P$ in the PI algorithm. We note that the authors in [10], [47], and [48] also propose the combination of Anderson mixing with optimal control algorithms. However, the QPI algorithm is fundamentally different in the sense that it approximates the transition matrix P using a different set of constraints that are specific to the optimal control algorithms.

B. QPI Algorithm

For $k \in \{0, 1, 2, \dots\}$, let

$$c_k := c^{\pi_{v_k}}, \quad P_k := P^{\pi_{v_k}}, \quad T_k := T(v_k).$$

(Recall that $c^{\pi_{v_k}}$ and $P^{\pi_{v_k}}$ are the stage cost and the state transition matrix of the greedy policy π_{v_k} w.r.t. v_k , respectively.) Recall the PI update rule

$$v_{k+1} = (I - \gamma P_k)^{-1} c_k = v_k - (I - \gamma P_k)^{-1} (v_k - T_k).$$

Inspired by the QNM approximation scheme (6), we propose the generic QPI update rule

$$v_{k+1} = v_k - (I - \gamma \tilde{P}_k)^{-1} (v_k - T_k) \quad (8)$$

where

$$\begin{aligned} \tilde{P}_k &= \arg \min_{P \in \mathbb{R}^{n \times n}} \|P - P_{\text{prior}}\|_F^2 \\ \text{s.t. } P r_i &= b_i, \quad i = 1, \dots, j. \end{aligned} \quad (9)$$

Observe that instead of approximating the complete Hessian $H_k := (I - \gamma P_k)^{-1}$ similar to standard QNMs, we are only approximating P_k . This choice particularly allows us to exploit the problem structure in order to form novel constraints and prior as we discuss next.

Regarding the constraints, the problem structure gives us two linear equality constraints: First, P_k is a row stochastic matrix, i.e., we have the *global* structural constraint

$$P_k \mathbf{1} = \mathbf{1} \quad (10)$$

and hence, we can set $r_1 = b_1 = \mathbf{1}$. Second, we can use the fact that the Bellman operator T is piecewise affine. In particular, from the definition (3) of the Bellman operator, it follows that $T(v) = c^{\pi_v} + \gamma P^{\pi_v} v$. Thus, we have the *local* structural constraint

$$T_k = c_k + \gamma P_k v_k \Rightarrow P_k v_k = \gamma^{-1} (T_k - c_k) \quad (11)$$

and we can set $r_2 = v_k$ and $b_2 = \gamma^{-1} (T_k - c_k)$. Note that, unlike the standard secant conditions in QNMs, the constraints (10) and (11) hold *exactly*. Incorporating these constraints, we propose the approximation

$$\begin{aligned} \tilde{P}_k &= \arg \min_{P \in \mathbb{R}^{n \times n}} \|P - P_{\text{prior}}\|_F^2 \\ \text{s.t. } P \mathbf{1} &= \mathbf{1}, \quad P v_k = \gamma^{-1} (T_k - c_k). \end{aligned} \quad (12)$$

The update rule (8) using the approximation (12) is, however, not necessarily a contraction. The same problem also arises in similar algorithms, such as Anderson accelerated VI [10] and Nesterov accelerated VI [9]. Here, we follow the standard

solution for this problem, that is, safeguarding the QPI update against the standard VI update based on the Bellman error

$$\theta_k := \|v_k - T_k\|_{\infty}.$$

To be precise, at each iteration $k = 0, 1, \dots$, we consider the *safeguarded* QPI update rule as follows:

(QPI) compute v_{k+1} according to (8), (12)

(Safeguard) if $\theta_{k+1} > \gamma^{k+1} \theta_0$, then $v_{k+1} = T_k$. (13)

The following theorem summarizes the discussion above by providing the QPI update rule explicitly (see Section IV-A for the proof).

Theorem 3.1 (QPI): Consider the update rule (8) using the approximation (12) where $P_{\text{prior}} \mathbf{1} = \mathbf{1}$ and let $G_{\text{prior}} = (I - \gamma P_{\text{prior}})^{-1}$. We have

$$v_{k+1} = v_k - \tilde{G}_k (v_k - T_k) \quad (14a)$$

where

$$\begin{aligned} w_k &= T_k - c_k - \gamma P_{\text{prior}} v_k, \quad \tilde{w}_k = G_{\text{prior}} w_k \in \mathbb{R}^n \\ u_k &= v_k - \frac{\mathbf{1}^\top v_k}{n} \mathbf{1}, \quad \tilde{u}_k = G_{\text{prior}}^\top u_k \in \mathbb{R}^n \\ \tau_k &= \begin{cases} 0, & \text{if } u_k^\top v_k = 0 \\ (u_k^\top v_k)^{-1}, & \text{otherwise} \end{cases} \in \mathbb{R} \\ \eta_k &= \begin{cases} 0, & \text{if } u_k^\top v_k = 0 \\ (u_k^\top (v_k - \tilde{w}_k))^{-1}, & \text{otherwise} \end{cases} \in \mathbb{R} \\ \tilde{P}_k &= P_{\text{prior}} + \gamma^{-1} \tau_k w_k u_k^\top \\ \tilde{G}_k &= G_{\text{prior}} + \eta_k \tilde{w}_k \tilde{u}_k^\top. \end{aligned} \quad (14b)$$

Moreover, the iterates v_k of the QPI update rule (14) with the safeguarding (13) converge to v^* linearly with rate γ and a per-iteration time complexity of $\mathcal{O}(n^2 m)$. $\square$

Observe that the safeguarded QPI update rule has the same per-iteration complexity as VI. Moreover, the convergence of QPI is ensured via safeguarding against VI, which leads to the same theoretically guaranteed linear convergence with rate γ as for VI. However, as we will show in the numerical examples below, we observe an empirically faster convergence for QPI with its rate showing less sensitivity to γ similar to QNMs. We also note that there is also a one-time computational cost of $\mathcal{O}(n^3)$ in the QPI update rule (14) for computing G_{prior} (assuming direct inversion and if G_{prior} is not available in closed form).

1) Prior: Next to be addressed is the choice of P_{prior} . First, note that for a fixed prior in all iterations, computation of τ_k and $\tilde{P}_k$ in (14b) is not needed since the update (14a) only requires $\tilde{G}_k$. The first choice for such a fixed prior is to exploit the available knowledge on the structure of the MDP. For instance, one can set $P_{\text{prior}} = P^\mu$ with μ being the stochastic policy choosing actions uniformly at random (so that P_{prior} is the average over actions of and has the same sparsity pattern as the true transition kernel of the MDP). A computationally advantageous choice for the prior is the uniform distribution $P_{\text{prior}} = \frac{1}{n} E = \frac{1}{n} \mathbf{1} \mathbf{1}^\top$ for which the update rule can be simplified significantly (see Section IV-B for the proof).

Corollary 3.2 (Uniform prior): The QPI update rule (14) with the uniform prior $P_{\text{prior}} = \frac{1}{n} E$, equivalently reads as

$$v_{k+1} = (1 - \delta_k) T_k + \delta_k c_k + \lambda_k \mathbf{1} \quad (15a)$$

where the scalar coefficients are given by

$$\begin{aligned} z_k &= c_k - \frac{\mathbf{1}^\top c_k}{n} \mathbf{1} \in \mathbb{R}^n \\ g_k &= v_k - T_k \in \mathbb{R}^n, \quad y_k = g_k - \frac{\mathbf{1}^\top g_k}{n} \mathbf{1} \in \mathbb{R}^n \\ \delta_k &= \begin{cases} 0, & \text{if } v_k^\top (y_k + z_k) = 0 \\ \frac{v_k^\top y_k}{v_k^\top (y_k + z_k)}, & \text{otherwise} \end{cases} \in \mathbb{R} \\ \lambda_k &= \frac{\gamma}{n(1-\gamma)} \mathbf{1}^\top ((\delta_k - 1)g_k + \delta_k c_k) \in \mathbb{R}. \end{aligned} \quad (15b)$$

□

Observe that the QPI update rule (15a) is a modification of the standard VI update rule $v_{k+1} = T(v_k)$ using two new vectors, namely, $c_k - T(v_k)$ and the all-one vector $\mathbf{1}$, with adaptive coefficients δ_k and λ_k , respectively.

Another interesting choice for the prior in (12) is $P_{\text{prior}} = \tilde{P}_{k-1}$, i.e., the previous approximation. This leads to a recursive scheme for approximating the transition matrix similar to QNM with Broyden approximation [46]. This can be achieved by choosing an initialization $\tilde{P}_{-1}$ such that $\tilde{P}_{-1}\mathbf{1} = \mathbf{1}$ and defining $\tilde{G}_{-1} := (1 - \gamma\tilde{P}_{-1})^{-1}$ (e.g., $\tilde{P}_{-1} = \frac{1}{n}E$ and $\tilde{G}_{-1} = I + \frac{\gamma}{n(1-\gamma)}E$). We note that for this recursive scheme, one might also need to introduce some form of projection to avoid unbounded growth of $\tilde{P}_k$ and $\tilde{G}_k$. In particular, using the fact that $\tilde{P}_k$ is in principle an approximation of a row stochastic matrix, one can use a simple normalization to ensure $\|\tilde{P}_k\|_\infty \leq 1$ and $\|\tilde{G}_k\|_\infty \leq (1 - \gamma)^{-1}$ for all k .

2) Implementation via Backtracking: The QPI update rule (14a) can be alternatively implemented without safeguarding. The proposed alternative again exploits the fact that the VI update leads to contraction in the Bellman error and combines that with the *backtracking* approach from optimization. To be precise, let us rewrite the QPI update rule (14a) as

$$\begin{aligned} v_{k+1} &= v_k - (\tilde{G}_k + I - I)(v_k - T_k) \\ &= T_k - (\tilde{G}_k - I)(v_k - T_k) \end{aligned}$$

and introduce a step-size α_k as follows:

$$v_{k+1} = T_k - \alpha_k(\tilde{G}_k - I)(v_k - T_k). \quad (16)$$

Then, at each iteration k , we introduce an inner iteration that finds the step-size α_k via backtracking such that $\theta_{k+1} \leq \gamma'\theta_k$ for a chosen $\gamma' \in (\gamma, 1)$. For instance, at each iteration k , we can set

- (0) $\alpha_k = 1$
- (1) compute v_{k+1} according to (14b) and (16)
- (2) if $\theta_{k+1} > \gamma'\theta_k$, then $\alpha_k \leftarrow \frac{1}{2}\alpha_k$, go to (1). (17)

We note that the preceding backtracking scheme is guaranteed to terminate if $v_k \neq v^*$ (i.e., $\theta_k \neq 0$). Indeed, we have the following (see Section IV-C for the proof).

Lemma 3.3 (Backtracking): Let $\theta_v := \|v - T(v)\|_\infty$ for $v \in \mathbb{R}^n$. Fix $v \in \mathbb{R}^n$ such that $\theta_v \neq 0$. Then, for all $w \in \mathbb{R}^n$ with $\|w - T(v)\|_\infty \leq \frac{\gamma'-\gamma}{1+\gamma}\theta_v$, we have $\theta_w \leq \gamma'\theta_v$. □

In particular, the preceding lemma implies that the backtracking scheme (17) terminates in at most $\log_2(\frac{1+\gamma}{\gamma'-\gamma}\|\tilde{G}_k - I\|_\infty)$ steps at iteration k . We also note that iterates of the proposed QPI update rule with backtracking converge to the optimal value function linearly with rate γ' .

3) Modified Implementation With Superlinear Convergence: For the proposed QPI scheme to achieve the classic local *superlinear* convergence of the class of QNMs, we need to modify the approximation (12) to also include secant-type conditions. To this end, for $k \geq 1$, we consider a modified approximation as follows:

if $\pi_{v_k} = \pi_{v_{k-1}}$

$$\begin{aligned} \tilde{P}_k &= \arg \min_{\tilde{P} \in \mathbb{R}^{n \times n}} \|P - \tilde{P}_{k-1}\|_F^2 \\ \text{s.t. } & P\mathbf{1} = \mathbf{1} \\ & Pv_k = \gamma^{-1}(T_k - c_k) \\ & \gamma P(v_k - v_{k-1}) = T_k - T_{k-1} \end{aligned} \quad (18a)$$

else

$$\begin{aligned} \tilde{P}_k &= \arg \min_{\tilde{P} \in \mathbb{R}^{n \times n}} \|P - \tilde{P}_{k-1}\|_F^2 \\ \text{s.t. } & P\mathbf{1} = \mathbf{1} \\ & \gamma P(v_k - v_{k-1}) = T_k - T_{k-1} \end{aligned} \quad (18b)$$

with the corresponding gain matrix

$$\tilde{G}_k = (I - \gamma\tilde{P}_k)^{-1}. \quad (18c)$$

Above, the only difference between (18a) and (18b) is the omission of the *local structural* constraint in (18b). We note that the condition $\pi_{v_k} = \pi_{v_{k-1}}$ implies that v_k and v_{k-1} have the same greedy policy, and hence lie on the same affine region of the map T (recall that T is piecewise affine). Therefore, the selection rule is to avoid “constraint collision” between the *local structural* and the *secant* constraints when $\pi_{v_k} \neq \pi_{v_{k-1}}$. Also, observe that in both approximations, the prior is the previous approximation $\tilde{P}_{k-1}$. Most importantly, note that the solution $\tilde{P}_k$ is a low-rank (rank-three, at most) update of $\tilde{P}_{k-1}$, cf., the optimization (6) and its solution (7). This means that $\tilde{G}_k$ can also be computed with a low cost using the Woodbury formula and $\tilde{G}_{k-1} = (I - \gamma\tilde{P}_{k-1})^{-1}$. Our next result concerns the convergence of the proposed modified implementation of the QPI algorithm (see Section IV-D for the proof).

Theorem 3.4 (Superlinear convergence): Consider the update rule (16) with the gain matrix (18) and the backtracking scheme (17). Then, the iterates v_k converge to v^* with a global linear rate and a local superlinear rate. □

4) Other Constraints: We finish this section with the following remark. One can also add extra constraints to the minimization problem (12) to impose a particular structure on the approximate transition matrix $\tilde{P}_k$. For instance, a natural constraint is to require this matrix to be entrywise nonnegative so that $\tilde{P}_k$ is indeed a probability transition matrix; or, one can impose a sparsity pattern on $\tilde{P}_k$ using the existing knowledge on the structure of the MDP. However, incorporating such information may lead to the problem (12) not having a closed-form and/or low-rank solution, and hence undermining the computational efficiency of the proposed algorithm. In this regard, we note that the problem (12) has a closed-form solution, which is a rank-one update of the prior; see $\tilde{P}_k$ in (14b).

C. Extension to Model-Free Control: QPL Algorithm

We now introduce the quasi-QPL algorithm as the extension of QPI for model-free control problems with access to samples

through a generative model. For simplicity, we limit the following discussion to the extension of the QPI algorithm (15) with a uniform prior. However, we note that the extension can be similarly applied to the QPI algorithm (14) with the generic prior.

The basic idea is to implement the stochastic version of the QPI update rule for the Q -function using the samples. In particular, similar to the approximation (12), we use an approximation of the state-action transition matrix under the greedy policy w.r.t. the Q -function q_k at each iteration k . However, the *local* structural constraint is in this case formed based on the *sampled* Bellman operator $\hat{T}(\cdot)$, evaluated at the sampled next states $\hat{s}_k^+$ at iteration k , as a surrogate for the Bellman operator $T(\cdot)$. To be precise, let

$$\hat{T}_k := \hat{T}(q_k)$$

at each iteration k . Also, let $c \in \mathbb{R}^{nm}$ be the vector of stage cost (with the same state-action ordering as the Q -function $q_k \in \mathbb{R}^{nm}$). We note that since the proposed QPL algorithm is synchronous with one sample for each state-action pair in each iteration, we have access to the complete stage cost c after the first iteration and can treat it as an input to the algorithm. The approximate state-action transition matrix $\tilde{P}_k$ at each iteration k is then formed as follows:

$$\begin{aligned} \tilde{P}_k &= \arg \min_{P \in \mathbb{R}^{nm \times nm}} \|P - P_{\text{prior}}\|_F^2 \\ \text{s.t. } P\mathbf{1} &= \mathbf{1}, Pq_k = \gamma^{-1}(\hat{T}_k - c). \end{aligned} \quad (19)$$

The minimization problem above also has a closed-form solution as a rank-one update of the prior, which allows us to compute $\tilde{G}_k = (I - \gamma\tilde{P}_k)^{-1}$ efficiently using the Woodbury formula. Having $\tilde{G}_k$ at hand, we can consider the update rule

$$q_{k+1} = q_k - \alpha_k \tilde{G}_k(q_k - \hat{T}_k) \quad (20)$$

where α_k is the diminishing learning rate of the algorithm, e.g., $\alpha_k = 1/(k+1)$. However, similar to the model-based case, the update rule (20) is not convergent. To address this issue, we consider a convex combination of the update rule (20) with the “standard” synchronous QL update rule³ [49], [50]

$$q_{k+1} = q_k - \alpha_k(q_k - \hat{T}_k) \quad (21)$$

with a diminishing effect. That is, we consider the update rule

$$\begin{aligned} q_{k+1} &= q_k - \alpha_k \left((1 - \beta_k)(q_k - \hat{T}_k) + \beta_k \tilde{G}_k(q_k - \hat{T}_k) \right) \\ &= q_k + \alpha_k(\hat{T}_k - q_k) + \alpha_k \beta_k (\tilde{G}_k - I)(\hat{T}_k - q_k) \end{aligned}$$

where β_k is a diminishing coefficient, e.g., $\beta_k = 1/(k+1)$. Moreover, to assure convergence, we bound the effect of the extra term $p_k = (\tilde{G}_k - I)(\hat{T}_k - q_k)$ by passing it through the operator $\Pi_M : \mathbb{R}^{nm} \rightarrow \mathbb{R}^{nm}$, with $M = \frac{2\gamma}{(1-\gamma)^2} \|c\|_\infty$, given by

$$\Pi_M(p) := \frac{\min\{M, \|p\|_\infty\}}{\|p\|_\infty} p.$$

³This is the so-called *synchronous* update of the Q -function in *all* state-action pairs in each iteration, corresponding to the *parallel sampling model* introduced by Kearns and Singh [49].

We note that the bound M is chosen using the fact that for any row stochastic matrix P and $G = (1 - \gamma P)^{-1}$, we have

$$\|G - I\|_\infty = \left\| \sum_{i=1}^{\infty} (\gamma P)^i \right\|_\infty \leq \sum_{i=1}^{\infty} \gamma^i \|P^i\|_\infty \leq \frac{\gamma}{1-\gamma}$$

and the generic bound $\|q^\pi\|_\infty \leq \|c\|_\infty / (1 - \gamma)$ for the value q^π of any policy π . The update rule of the model-free QPL algorithm in its generic form is hence

$$\begin{aligned} p_k &= (\tilde{G}_k - I)(\hat{T}_k - q_k) \\ q_{k+1} &= q_k + \alpha_k(\hat{T}_k - q_k) + \alpha_k \beta_k \Pi_M(p_k). \end{aligned} \quad (22)$$

In particular, by using the uniform prior $P_{\text{prior}} = \frac{1}{nm}E$, the QPL update rule reduces to

$$\begin{aligned} p_k &= \delta_k(c - \hat{T}_k) + \lambda_k \mathbf{1} \\ q_{k+1} &= q_k + \alpha_k(\hat{T}_k - q_k) + \alpha_k \beta_k \Pi_M(p_k) \end{aligned} \quad (23a)$$

where

$$\begin{aligned} z &= c - \frac{\mathbf{1}^\top c}{nm} \mathbf{1} \in \mathbb{R}^{nm} \\ \hat{g}_k &= q_k - \hat{T}_k \in \mathbb{R}^{nm}, \quad y_k = \hat{g}_k - \frac{\mathbf{1}^\top \hat{g}_k}{nm} \mathbf{1} \in \mathbb{R}^{nm} \\ \delta_k &= \begin{cases} 0, & \text{if } q_k^\top(y_k + z) = 0 \\ \frac{q_k^\top y_k}{q_k^\top(y_k + z)}, & \text{otherwise} \end{cases} \in \mathbb{R} \\ \lambda_k &= \frac{\gamma}{nm(1-\gamma)} \mathbf{1}^\top ((\delta_k - 1)\hat{g}_k + \delta_k c) \in \mathbb{R}. \end{aligned} \quad (23b)$$

Observe that in this case, QPL uses the two additional vectors $c - \hat{T}_k$ and the all-one vector $\mathbf{1}$ with adaptive coefficients in its update rule. The following theorem summarizes properties of the proposed QPL algorithm (see Section IV-E for the proof).

Theorem 3.5 (QPL): Assume the learning rates α_k and β_k are such that $\sum_k \alpha_k = \infty$, $\sum_k \alpha_k^2 < \infty$, and $\beta_k \rightarrow 0$. Then, the iterates q_k of QPL algorithm (23) converge to q^* almost surely. Moreover, the algorithm has a per-iteration complexity of $\mathcal{O}(nm^2)$. $\square$

Regarding the preceding result, we note that the per-iteration time complexity of QPL is the same as that of the (synchronous) QL algorithm. We finish with the following remark on the *asynchronous* implementation of QPL.

Remark 3.6 (Asynchronous QPL): The proposed QPL update rule (23) can also be implemented in an asynchronous fashion. To be precise, this requires forming the approximate state-action transition matrix $\tilde{P}_k$ based on a single sample $(s_k, a_k, \hat{s}_k^+, c(s_k, a_k))$ as follows:

$$\begin{aligned} \tilde{P}_k &= \arg \min_{P \in \mathbb{R}^{nm \times nm}} \|P - P_{\text{prior}}\|_F^2 \\ \text{s.t. } P\mathbf{1} &= \mathbf{1}, \end{aligned}$$

$$e_{(s_k, a_k)}^\top P q_k = \gamma^{-1}(\hat{T}_k(s_k, a_k) - c(s_k, a_k))$$

cf., approximation (19). In particular, by using the uniform prior $P_{\text{prior}} = \frac{1}{nm}E$, the corresponding update rule reads as

$$\begin{aligned} p_k &= \tau_k \left(\delta_k e_{(s_k, a_k)} + \frac{\gamma(1 + \delta_k)}{nm(1 - \gamma)} \mathbf{1} \right) \\ q_{k+1} &= q_k + \alpha_k \tau_k e_{(s_k, a_k)} + \alpha_k \beta_k \Pi_M(p_k) \end{aligned}$$

where the scalar coefficients are given by

$$\tau_k = \hat{T}_k(s_k, a_k) - q_k(s_k, a_k), \quad \rho_k = \frac{1}{nm} \mathbf{1}^\top q_k$$

$$\lambda_k = \left(\widehat{T}_k(s_k, a_k) - c(s_k, a_k) - \gamma \rho_k \right) (q_k(s_k, a_k) - \rho_k)$$

$$\eta_k = \|q_k\|_2^2 - \rho_k^2 - \lambda_k, \quad \delta_k = \begin{cases} 0, & \text{if } \eta_k = 0 \\ \lambda_k / \eta_k, & \text{otherwise.} \end{cases}$$

Note that the preceding update rule leads to an update in all entries of the Q-function q_k in each iteration. $\square$

IV. TECHNICAL PROOFS

A. Proof of Theorem 3.1

First, let us show that the two equality constraints in the minimization problem (12) are linearly dependent if and only if $u_k^\top v_k = 0$. In this regard, observe that the two equality constraints are linearly dependent if and only if $v_k = \rho \mathbf{1}$ for some $\rho \in \mathbb{R}$: For $\rho = 0$, the second equality constraint becomes trivial; and, for $\rho \neq 0$, the two constraints become equivalent. On the other hand, we have

$$u_k^\top v_k = \left(v_k - \frac{\mathbf{1}^\top v_k}{n} \mathbf{1} \right)^\top v_k = v_k^\top \left(I - \frac{1}{n} E \right) v_k.$$

Then, since $I - \frac{1}{n} E$ is positive semidefinite with one zero eigenvalue corresponding to the eigenvector $\mathbf{1}$, we have $u_k^\top v_k = 0$ if and only if $v_k = \rho \mathbf{1}$ for some $\rho \in \mathbb{R}$. Hence, the constraints in (12) are linearly dependent if and only if $u_k^\top v_k = 0$.

We first consider the update rule (14) for the case $u_k^\top v_k \neq 0$. Define $R := (\mathbf{1}, v_k)$ and $B := (\mathbf{1}, \gamma^{-1}(T_k - c_k)) \in \mathbb{R}^{n \times 2}$ so that the minimization problem (12) can be written as

$$\tilde{P}_k = \operatorname{argmin}_{P \in \mathbb{R}^{n \times n}} \{ \|P - P_{\text{prior}}\|_F^2 : PR = B \}.$$

Note that, since $u_k^\top v_k \neq 0$ and hence v_k and $\mathbf{1}$ are linearly independent, R is of full column rank. The solution to the preceding problem is given by

$$\tilde{P}_k = P_{\text{prior}} + (B - P_{\text{prior}}R)(R^\top R)^{-1}R^\top.$$

Now, observe that

$$\begin{aligned} (B - P_{\text{prior}}R) &= \begin{bmatrix} \mathbf{1} & \frac{1}{\gamma}(T_k - c_k) \end{bmatrix} - P_{\text{prior}} \begin{bmatrix} \mathbf{1} & v_k \end{bmatrix} \\ &= \begin{bmatrix} \mathbf{1} - P_{\text{prior}}\mathbf{1} & \frac{1}{\gamma}(T_k - c_k) - P_{\text{prior}}v_k \end{bmatrix} \\ &= \begin{bmatrix} \mathbf{0} & \frac{1}{\gamma}(T_k - c_k) - P_{\text{prior}}v_k \end{bmatrix} \\ &= \begin{bmatrix} \mathbf{0} & \gamma^{-1}w_k \end{bmatrix} \end{aligned}$$

where we used the assumption $P_{\text{prior}}\mathbf{1} = \mathbf{1}$. Also

$$\begin{aligned} (R^\top R)^{-1} &= \left(\begin{bmatrix} \mathbf{1}^\top \\ v_k^\top \end{bmatrix} \begin{bmatrix} \mathbf{1} & v_k \end{bmatrix} \right)^{-1} = \begin{bmatrix} n & v_k^\top \mathbf{1} \\ v_k^\top \mathbf{1} & v_k^\top v_k \end{bmatrix}^{-1} \\ &= \frac{1}{n(u_k^\top v_k)} \begin{bmatrix} v_k^\top v_k & -\mathbf{1}^\top v_k \\ -\mathbf{1}^\top v_k & n \end{bmatrix} \end{aligned}$$

and

$$(R^\top R)^{-1}R^\top = \frac{1}{u_k^\top v_k} \begin{bmatrix} * \\ u_k^\top \end{bmatrix}.$$

Therefore, we have

$$\tilde{P}_k = P_{\text{prior}} + \gamma^{-1}(u_k^\top v_k)^{-1}w_k u_k^\top.$$

For the case $u_k^\top v_k = 0$, as we discussed in the beginning of the proof, the two equality constraints in the minimization problem (12) become linearly dependent, and, in particular, the second constraint can be discarded. The solution to the problem (12)

in this case is then $\tilde{P}_k = P_{\text{prior}}$. Hence, the approximation (12) can be in general written as

$$\tilde{P}_k = P_{\text{prior}} + \gamma^{-1}\tau_k w_k u_k^\top \quad (24)$$

where

$$\tau_k = \begin{cases} 0, & \text{if } u_k^\top v_k = 0 \\ (u_k^\top v_k)^{-1}, & \text{otherwise.} \end{cases}$$

That is, the approximation $\tilde{P}_k$ is a rank-one update of the prior. Then, if $\tau_k = 0$, we clearly have

$$\tilde{G}_k = (I - \gamma\tilde{P}_k)^{-1} = (I - \gamma P_{\text{prior}})^{-1} = G_{\text{prior}}$$

and, for the case $\tau_k \neq 0$, we can use the Woodbury formula to write

$$\begin{aligned} \tilde{G}_k &= (G_{\text{prior}}^{-1} - \tau_k w_k u_k^\top)^{-1} \\ &= G_{\text{prior}} + \frac{1}{\tau_k^{-1} - u_k^\top (G_{\text{prior}} w_k)} (G_{\text{prior}} w_k) (G_{\text{prior}}^\top u_k)^\top \\ &= G_{\text{prior}} + \eta_k \check{w}_k \check{u}_k^\top. \end{aligned}$$

What remains to be shown is that η_k is well-defined for $u_k^\top v_k \neq 0$ (i.e., $\tau_k \neq 0$). First, we use $T_k = c_k + \gamma P_k v_k$ to write

$$\begin{aligned} \eta_k^{-1} &= u_k^\top (v_k - \check{w}_k) \\ &= u_k^\top (v_k - G_{\text{prior}}(T_k - c_k - \gamma P_{\text{prior}} v_k)) \\ &= u_k^\top (v_k - G_{\text{prior}}(\gamma P_k v_k - \gamma P_{\text{prior}} v_k)) \end{aligned}$$

Next, since $P_{\text{prior}} = \gamma^{-1}(I - G_{\text{prior}}^{-1})$ and using the fact that $v_k = u_k + \frac{\mathbf{1}^\top v_k}{n} \mathbf{1}$, we have

$$\begin{aligned} \eta_k^{-1} &= u_k^\top (I - \gamma G_{\text{prior}}(P_k - P_{\text{prior}})) v_k \\ &= u_k^\top G_{\text{prior}}(I - \gamma P_k) v_k \\ &= u_k^\top G_{\text{prior}}(I - \gamma P_k)(u_k + \alpha_k \mathbf{1}) \\ &= u_k^\top G_{\text{prior}}(I - \gamma P_k) u_k + \frac{\mathbf{1}^\top v_k}{n} u_k^\top G_{\text{prior}}(I - \gamma P_k) \mathbf{1}. \end{aligned}$$

Now, note that $P_k \mathbf{1} = \mathbf{1}$ and $G_{\text{prior}} \mathbf{1} = (1 - \gamma)^{-1} \mathbf{1}$. Hence

$$\begin{aligned} \eta_k^{-1} &= u_k^\top G_{\text{prior}}(I - \gamma P_k) u_k + \frac{\mathbf{1}^\top v_k}{n} u_k^\top \mathbf{1} \\ &= u_k^\top G_{\text{prior}}(I - \gamma P_k) u_k \end{aligned}$$

where we also used the fact that $u_k^\top \mathbf{1} = 0$. Then, since the matrices G_{prior} and $(I - \gamma P_k)$ are nonsingular with their eigenvalues having strictly positive real parts (because the eigenvalues of P_k all reside within the unit disc), we have

$$\begin{aligned} u_k^\top G_{\text{prior}}(I - \gamma P_k) u_k &= 0 \Leftrightarrow u_k = \mathbf{0} \\ &\Leftrightarrow v_k = \rho \mathbf{1} \text{ for some } \rho \in \mathbb{R} \\ &\Leftrightarrow u_k^\top v_k = 0. \end{aligned}$$

That is, η_k is well-defined for $u_k^\top v_k \neq 0$.

Next, we consider the rate of convergence of the safeguarded QPI update rule. We will use induction to show that

$$\theta_k \leq \gamma^k \theta_0 \quad \forall k \geq 0. \quad (25)$$

The base case $k = 0$ holds trivially. Let us now assume the inequality (25) holds for some $k \geq 0$ and recall that T is a γ -contraction in ∞ -norm. Then, the safeguarding against standard

VI as in (13) implies that

$$\theta_{k+1} \leq \max\{\gamma^{k+1}\theta_0, \gamma\theta_k\} \leq \gamma^{k+1}\theta_0$$

where we used the induction hypothesis for the last inequality. This completes the proof.

Finally, the per-iteration time complexity of each iteration of the safeguarded QPI update rule: The update rule (14a) requires $\mathcal{O}(n^2 m)$ operations for computing the vectors $T_k = T(v_k)$, $\mathcal{O}(n^2)$ operations for the matrix-vector multiplication, and $\mathcal{O}(n)$ operations for the vector additions. Computing the objects in (14b) involves vector/matrix additions and matrix-vector multiplications (all of size n), and hence requires $\mathcal{O}(n^2)$ operations. For the safeguarding (13), we need to compute $T_{k+1} = T(v_{k+1})$, which again requires $\mathcal{O}(n^2 m)$ operations. Summing up the aforementioned complexities, we derive the total time complexity to be $\mathcal{O}(n^2 m)$.

B. Proof of Corollary 3.2

The result follows from Theorem 3.1 by plugging in $P_{\text{prior}} = \frac{1}{n}E$ and $G_{\text{prior}} = I + \frac{\gamma}{n(1-\gamma)}E$ and simplifying the expression. In particular, we note that the condition $v_k^\top(y_k + z_k) = 0$ in (15b) is equivalent to the condition $u_k^\top v_k = 0$ in (14b). To see this, recall that $u_k^\top v_k = 0$ if and only if $v_k = \rho \mathbf{1}$ for some $\rho \in \mathbb{R}$; see the first part of the proof of Theorem 3.1 in Section IV-A. Also, observe that

$$\begin{aligned} v_k^\top(y_k + z_k) &= v_k^\top \left(g_k + c_k - \frac{\mathbf{1}^\top(g_k + c_k)}{n} \mathbf{1} \right) \\ &= v_k^\top \left(v_k - T_k + c_k - \frac{\mathbf{1}^\top(v_k - T_k + c_k)}{n} \mathbf{1} \right) \\ &= v_k^\top \left(I - \frac{1}{n}E \right) (v_k - T_k + c_k) \\ &= v_k^\top \left(I - \frac{1}{n}E \right) (I - \gamma P_k) v_k \end{aligned}$$

where, for the last equality, we used $T_k = c_k + \gamma P_k v_k$. Then, since $u_k = (I - \frac{1}{n}E)v_k = v_k - \frac{\mathbf{1}^\top v_k}{n} \mathbf{1}$, we have

$$\begin{aligned} v_k^\top(y_k + z_k) &= u_k^\top (I - \gamma P_k) \left(u_k + \frac{\mathbf{1}^\top v_k}{n} \mathbf{1} \right) \\ &= u_k^\top (I - \gamma P_k) u_k + \frac{\mathbf{1}^\top v_k}{n} u_k^\top (I - \gamma P_k) \mathbf{1} \\ &= u_k^\top (I - \gamma P_k) u_k + \frac{\mathbf{1}^\top v_k}{n} (1 - \gamma) u_k^\top \mathbf{1} \\ &= u_k^\top (I - \gamma P_k) u_k \end{aligned}$$

where, for the last equality, we used the fact that $u_k^\top \mathbf{1} = 0$. Finally, since $(I - \gamma P_k)$ is nonsingular with its eigenvalues having strictly positive real parts (because the eigenvalues of P_k all reside within the unit disc), we have

$$\begin{aligned} v_k^\top(y_k + z_k) &= 0 \Leftrightarrow u_k = \mathbf{0} \\ &\Leftrightarrow v_k = \rho \mathbf{1} \text{ for some } \rho \in \mathbb{R}. \end{aligned}$$

This completes the proof.

C. Proof of Lemma 3.3

Recall that T is a γ -contraction observe that

$$\theta_{T(v)} = \|T(v) - T(T(v))\|_\infty \leq \gamma \|v - T(v)\|_\infty = \gamma \theta_v.$$

Also, note that the map $\theta_{(\cdot)}$ is $(1 + \gamma)$ -Lipschitz continuous: For any $w_1, w_2 \in \mathbb{R}^n$, we have

$$\begin{aligned} |\theta_{w_1} - \theta_{w_2}| &= | \|w_1 - T(w_1)\|_\infty - \|w_2 - T(w_2)\|_\infty | \\ &\leq \|w_1 - T(w_1) - w_2 + T(w_2)\|_\infty \\ &\leq \|w_1 - w_2\|_\infty + \|T(w_1) - T(w_2)\|_\infty \\ &\leq (1 + \gamma) \|w_1 - w_2\|_\infty. \end{aligned}$$

Setting $w_1 = w$ with $\|w - T(v)\|_\infty \leq \frac{\gamma' - \gamma}{1 + \gamma} \theta_v$ and $w_2 = T(v)$, we then have

$$|\theta_w - \theta_{T(v)}| \leq (1 + \gamma) \|w - T(v)\|_\infty \leq (\gamma' - \gamma) \theta_v.$$

Hence

$$\theta_w \leq \theta_{T(v)} + (\gamma' - \gamma) \theta_v \leq \gamma' \theta_v.$$

This completes the proof.

D. Proof of Theorem 3.4

In what follows, we show that the proposed algorithm is an instance of semismooth QNM with a structured least-change secant update and a standard backtracking scheme. To this end, observe that we are looking for the unique root v^* of the map $g(v) = v - T(v)$. First, note that g is piecewise affine, Lipschitz-continuous, and, in particular, *strongly semismooth* [51, Prop. 7.4.7]. Moreover, any generalized Jacobian $H \in \partial g(v) \subset \{I - \gamma P : P \geq 0, P\mathbf{1} = \mathbf{1}\}$ is nonsingular at any point v (note that for a row stochastic matrix P and $\gamma \in (0, 1)$, we have $(I - \gamma P)^{-1} = \sum_{t=0}^{\infty} \gamma^t P^t$).

Denote $g_k := g(v_k)$ and $\tilde{H}_k := \tilde{G}_k^{-1}$ for each $k \geq 0$. The update rule (16) then can be equivalently written as

$$v_{k+1} = v_k - \alpha_k \tilde{H}_k^{-1} g_k - (1 - \alpha_k) g_k$$

where α_k is determined using the backtracking scheme (17) ensuring

$$\|g_{k+1}\|_\infty \leq \gamma' \|g_k\|_\infty$$

for some $\gamma' \in (\gamma, 1)$. This means that, as long as the approximations $\tilde{H}_k$ of the Jacobian of g are uniformly bounded, we have the global linear convergence of the iterates v_k to v^* with rate γ' ; see Lemma 3.3.

Finally, let us look at the approximations $\tilde{H}_k$ of the Jacobian of g . Using the change of variable $\tilde{P}_k = \gamma^{-1}(I - \tilde{H}_k)$, we can show that

$$\begin{aligned} \text{if } \pi_{v_k} &= \pi_{v_{k-1}} \\ \tilde{H}_k &= \arg \min_{H \in \mathbb{R}^{n \times n}} \|H - \tilde{H}_{k-1}\|_F^2 \\ \text{s.t. } & H\mathbf{1} = (1 - \gamma)\mathbf{1} \\ & H v_k = g_k + c_k \\ & H(v_k - v_{k-1}) = g_k - g_{k-1} \end{aligned}$$

else

$$\begin{aligned} \tilde{H}_k &= \arg \min_{H \in \mathbb{R}^{n \times n}} \|H - \tilde{H}_{k-1}\|_F^2 \\ \text{s.t. } & H\mathbf{1} = (1 - \gamma)\mathbf{1} \\ & H(v_k - v_{k-1}) = g_k - g_{k-1}. \end{aligned}$$

One can now clearly observe that the approximation $\tilde{H}_k$ satisfies the *secant* condition, i.e.,

$$\tilde{H}_k(v_k - v_{k-1}) = g_k - g_{k-1}$$

besides the structural linear constraints that the true (generalized) Jacobian must satisfy. This implies that the proposed algorithm indeed involves a least-change in $\|\tilde{H}_k - \tilde{H}_{k-1}\|_F$ subject to secant and structural constraints. Now, since g is strongly semismooth, there exists $H_k \in \partial g(v_k)$ such that [51, Def. 7.4.2]

$$\|g_k - g_{k-1} - H_k(v_k - v_{k-1})\| = o(\|v_k - v_{k-1}\|).$$

Therefore

$$\|(\tilde{H}_k - H_k)(v_k - v_{k-1})\| = o(\|v_k - v_{k-1}\|).$$

That is, Dennis–Moré condition is satisfied and, by [52, Thm. 4.2], the convergence is locally superlinear.

E. Proof of Theorem 3.5

The per-iteration time complexity of each iteration of the QPL update rule is as follows: The update rule (23) requires $\mathcal{O}(nm^2)$ operations for computing the vectors $\hat{T}_k = \hat{T}(q_k)$, and $\mathcal{O}(nm)$ operations for computing the step-sizes δ_k and λ_k and the vector additions. Summing up the aforementioned complexities, the total time complexity is $\mathcal{O}(nm^2)$.

Regarding the convergence, define

$$F(q) := \bar{T}(q) - q, \quad V_k := \hat{T}_k - T(q), \quad d_k := \beta_k \Pi_M(p_k)$$

where $\bar{T}$ is the corresponding true Bellman operator for Q-functions, i.e., $\bar{T}(q) = \mathbb{E}[\hat{T}(q)]$ for each q . Then, the convergence of q_k to the unique equilibrium q^* of the map F follows from [53, Corollary 3.2]. In particular, observe that $d_k \rightarrow 0$ since $\beta_k \rightarrow 0$ and $\|\Pi_M(p_k)\|_\infty \leq M$.

V. NUMERICAL SIMULATIONS

We now compare the performance of the proposed algorithms with that of the standard existing algorithms for the optimal control of different MDPs. See Appendix A for a description of the considered MDPs. To this end, we first focus on the proposed algorithms with uniform priors corresponding to update rules (15) and (23) in Sections V-A and V-B, respectively. The results of numerical experiments with alternative priors are then reported in Section V-C.

A. Model-Based Algorithms

For model-based algorithms, we consider two MDPs: a randomly generated Garnet MDP [54] and the Healthcare MDP [9] with an *absorbing* state. The proposed QPI algorithm (15) is compared with the following algorithms: VI; VI with Nesterov acceleration (NVI) [9]; VI with Anderson acceleration (AVI) [47]; and policy iteration. For AVI, we use a memory of one leading to a rank-one update (of the identity matrix) for approximating the Hessian so that it is comparable with the

rank-one update of the uniform distribution for approximating the transition matrix in QPI. See Appendix B for the exact update rules of NVI and AVI. We note that since NVI and AVI are not guaranteed to converge, we safeguard them using VI (using the same safeguarding rule (13) used for QPI). All the algorithms are initialized by $v_0 = \mathbf{0}$ with termination condition $\|v_k - T(v_k)\|_\infty \leq 10^{-6}$. The results of the simulations are provided in Figs. 1 and 2.

In Fig. 1, VI, NVI, and AVI show a linear convergence with a rate depending on γ in both MDPs. In particular, as we increase γ from 0.9 to 0.999, we observe more than a tenfold increase in the number of iterations required for these algorithms to terminate. This is expected since these algorithms only use first-order information and their convergence rate is determined by γ .

Fig. 1 also shows that for both MDPs, PI converges with a quadratic rate in 3 to 5 iterations, independent of γ . Now, observe that QPI is the only algorithm showing a similar behavior as PI and terminating in approximately the same number of iterations, independent of γ , in both MDPs. Moreover, comparing the performance of QPI with AVI (its counterpart in the class of QNMs), we also see the importance of newly introduced linear constraints and prior in the approximation of the transition matrix. Moreover, observe that QPI's safeguard is activated for Healthcare MDP, as shown in Fig. 1(b) (one instance for $\gamma = 0.99$ and multiple instances for $\gamma = 0.999$). In this regard, we note that for QPI, we have observed that the activation of the safeguard is particularly due to the existence of *absorbing states* in the MDP as is the case for the Healthcare MDP.

Fig. 2 reports the corresponding running times of the algorithms. The reported running times are in line with convergence behaviors seen in Fig. 1 and the theoretical time complexity of these algorithms. In particular, PI and QPI are the only algorithms with running time less sensitive to γ for both of the considered MDPs. In this regard, we note that since the size of the MDPs considered in our numerical simulations is relatively small, PI is the fastest algorithm despite the fact that it requires a matrix inversion.

B. Model-Free Algorithms

For model-free algorithms, we also consider two MDPs: again a randomly generated Garnet MDP [54] and the Graph MDP [21]. The proposed QPL algorithm (23) with $\alpha_k = \frac{1}{(k+1)}$ and $\beta_k = \frac{1}{(k+1)^{0.1}}$ is compared with the following algorithms: Q-learning as in (21) with $\alpha_k = \frac{1}{(k+1)}$; SQL QL [14]; and ZQL [21]. See Appendix B for the exact update rules of SQL and ZQL. All the algorithms are initialized by $q_0 = \mathbf{0}_{nm}$ and terminated after $K = 10^4$ iterations with a synchronous sampling of all state-action pairs at each iteration. For each algorithm, we report the *average* of the Bellman error $\|q_k - T(q_k)\|_\infty$ over 20 runs of the algorithm. The results of the simulations are provided in Fig. 3 and Table I.

As can be seen in Fig. 3(a) and (b), the performance of QL and SQL (the first-order methods) deteriorates as γ increases for both MDPs. However, for these MDPs, ZQL (the second-order method that estimates the transition matrix by averaging over

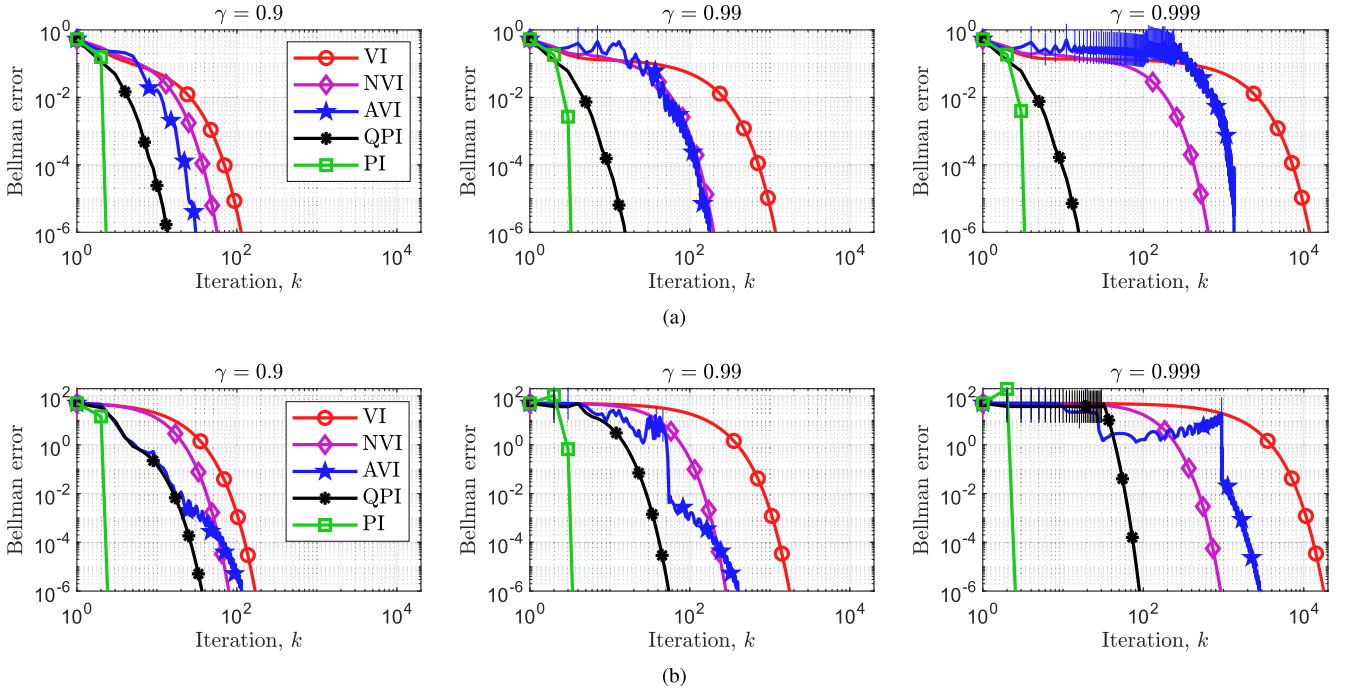


Fig. 1. Performance of model-based algorithms for three values of γ . (a) Garnet MDP. (b) Healthcare MDP. The bars indicate the iterations at which the safeguard is activated (for NVI, AVI, and QPI).

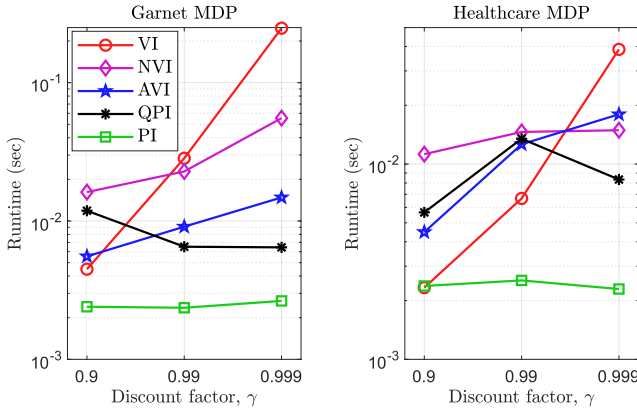


Fig. 2. Running time of the model-based algorithms for three values of γ corresponding to Fig. 1.

TABLE I
RUNNING TIME (IN SECONDS) OF THE MODEL-FREE ALGORITHMS OVER $K = 10^4$ ITERATIONS (AVERAGED OVER 20 RUNS) FOR $\gamma = 0.9$ CORRESPONDING TO FIG. 3

	QL	SQL	QPL	ZQL	Sampling
Garnet	1.7	3.1	3.3	16	71
Graph	0.16	0.27	0.31	0.59	6.5

the samples) leads to almost the same error level after a fixed number of iterations for different values of γ .

Fig. 3 shows that the performance of QPL is not as consistent as its model-based counterpart: QPL has the same rate of convergence as ZQL for Garnet MDP [see Fig. 3(a)], while it deviates

from ZQL as k increases for Graph MDP [see Figs. 3(b)]. This means that for structured MDPs, QPL may not lead to a better performance compared to SQL or ZQL. In this regard, we note that the model-free QPL algorithm uses an approximation of the transition matrix, which is constructed based on sampled data; see (19). This use of sampling on top of approximation can be the reason behind the poor performance of the model-free QPL algorithm for structured MDPs.

Finally, we note that the running times reported in Table I also align with the corresponding theoretical time complexities of these algorithms. In particular, QPL and SQL require almost the same amount of time, which is slightly more than QL and less than ZQL. (We report the runtime only for $\gamma = 0.9$ because it is independent of γ .) Note that the time required for generating the samples is reported separately in Table I, which is indeed the dominating factor in the actual runtime of the model-free algorithms.

C. QPI and QPL With Different Priors

Figs. 4 and 5 report the result of our numerical simulations for the QPI and QPL algorithms, respectively, with three choices of the prior: 1) QPI/L-A with a uniform prior $P_{\text{prior}} \propto \mathbf{1}\mathbf{1}^\top$, 2) QPI/L-B with recursive prior $P_{\text{prior}} = \tilde{P}_{k-1}$, and 3) QPI/L- μ with prior $P_{\text{prior}} = P^\mu$ and μ being the stochastic policy choosing actions uniformly at random so that the prior has the same sparsity pattern as the true transition probability matrix.

As depicted in Figs. 4(a) and 5(a), the experiments with alternative priors show no improvement in the performance of the QPI and QPL algorithms in comparison with the uniform prior for random Garnet MDPs. For structured MDPs, however,

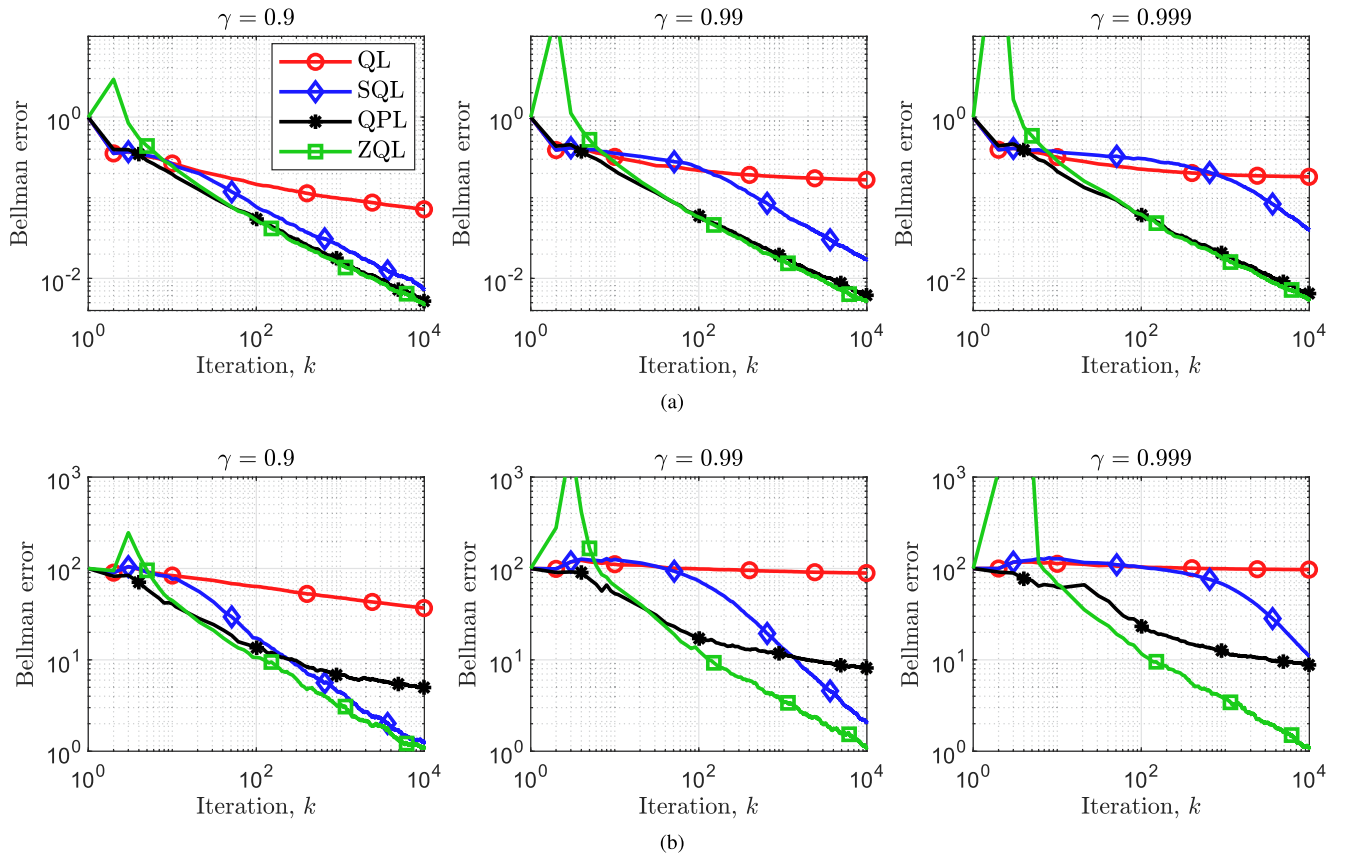


Fig. 3. Performance of model-free algorithms (averaged over 20 runs) for three values of γ . (a) Garnet MDP. (b) Graph MDP.

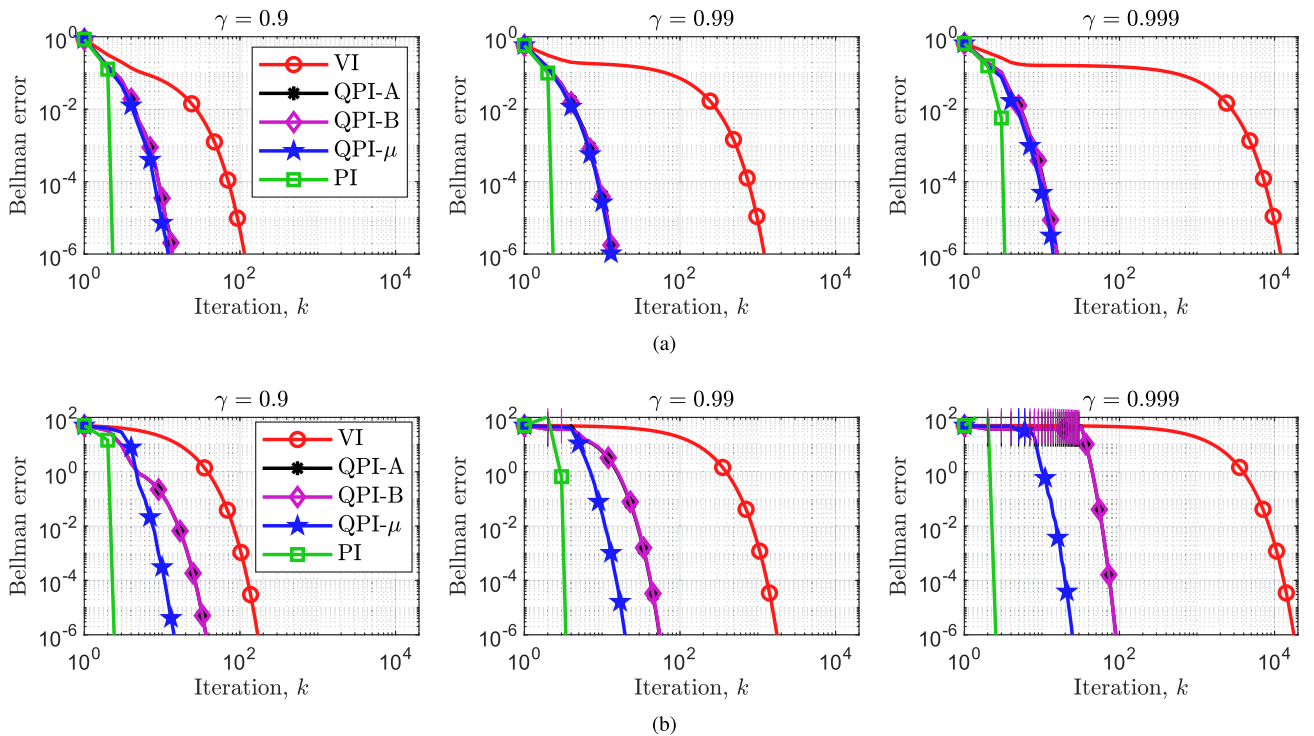


Fig. 4. Performance of model-based algorithms for three values of γ and three different priors. (a) Garnet MDP. (b) Healthcare MDP. The bars indicate the iterations at which the safeguard is activated in QPI.

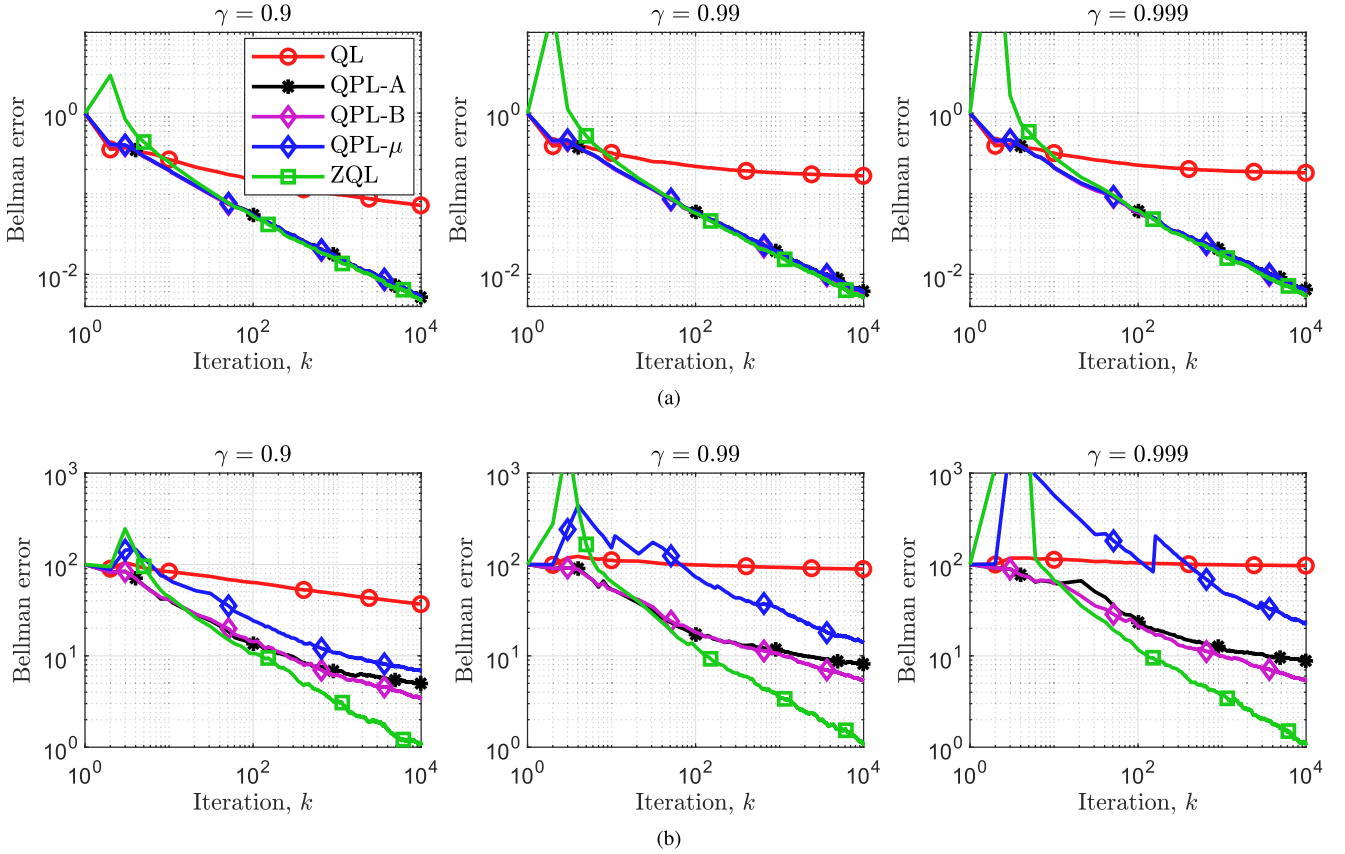


Fig. 5. Performance of model-free algorithms (averaged over 20 runs) for three values of γ and three different priors. (a) Garnet MDP. (b) Graph MDP.

we observe contradictory results, as shown in Figs. 4(b) and 5(b): Using a structured prior leads to a significant improvement in the performance of the (model-based) QPI algorithm for Healthcare MDP, while using a structured prior significantly deteriorates the performance of the (model-free) QPL algorithm for Graph MDP.

VI. LIMITATIONS AND FUTURE RESEARCH

In this article, we proposed the model-based QPI algorithm and its model-free counterpart, the QPL algorithm. The proposed algorithms were particularly inspired by the QNMs and employed a novel approximation of the “Hessian” by using two new linear constraints specific to MDPs.

The main drawback of the proposed algorithms, similar to other accelerated VI schemes in the literature, is the need for safeguarding to ensure convergence. Our experiments in Section V showed examples of MDPs in which the safeguard is activated. Our modified implementation of Section III-B2 guarantees the linear convergence while removing the need for safeguarding and using backtracking instead. Another possible solution is the use of the operator splitting method introduced in [55] for policy evaluation. In this regard, we note that the proposed QPI algorithm is essentially the PI algorithm in which the policy evaluation step uses the approximation $\tilde{P}_k$ in (12) instead of the true transition matrix P_k and the cost $\tilde{c}_k = c_k + \gamma(P_k - \tilde{P}_k)v_k$

instead of the true cost c_k . However, the convergence requires $\tilde{P}_k$ to be close to P_k . To be precise, a sufficient condition is $\|P_k - \tilde{P}_k\|_\infty \leq 1 - \gamma$ [55, Thm. 1], which is difficult to achieve for a low-rank approximation $\tilde{P}_k$ of P_k .

Another limitation of the current work is the lack of a theoretical guarantee for the empirically observed improvement in the convergence rate, particularly for the model-based QPI algorithm. The main difficulty to be addressed is the fact that the linear constraints in (12) are not the standard secant conditions used in QNMs. The modified implementation of Section III-B3 partially addresses this issue by including the secant conditions in its approximation. An interesting result, however, is the establishment of a local superlinear convergence rate for QPI without imposing the secant condition. Another possibility is to use the results for Anderson acceleration in [56] to establish an improved linear rate for convergence. To that end, similar to what is done in [48], one needs to use a smoothed version of the Bellman operator, e.g., by replacing the *max* operation with a *soft-max* operation in the Bellman operator.

The proposed algorithms in this study heavily rely on the approximation (12) of the transition matrix. As we discussed, this approximation easily allows for incorporation of different priors, e.g., a prior with the same sparsity pattern as the true transition matrix, or, the recursive prior. However, our numerical

simulations using these alternative priors did not demonstrate a clear improvement in the performance of the proposed algorithms, highlighting the need for further investigation on other MDPs. In this regard, we also note that the main drawback of the approximation (12) is that it does not allow for a computationally efficient incorporation of other constraints, such as nonnegativity constraints. A promising future research direction is the development of alternative approximation schemes that allow such constraints to be included at a reasonable computational cost.

APPENDIX

A. MDPs of the Numerical Simulations

Garnet MDP: The considered Garnet MDP [54] is generated randomly with $n = 50$ states, $m = 5$ actions, and the branching parameter $n_b = 10$. For each state-action pair (s, a) , we first form the set of reachable next states $\{s_1^+, \dots, s_{n_b}^+\}$ chosen uniformly at random from the state space $\{1, \dots, n\}$. Then, the corresponding probabilities are formed by choosing the points $p_i \in [0, 1]$, $i = 1, \dots, n_b - 1$, uniformly at random, and setting $\mathbb{P}(s_i^+ | s, a) = p_i - p_{i-1}$ with $p_0 = 0$ and $p_{n_b} = 1$. The stage cost $c(s, a)$ for each state-action pair (s, a) is also chosen uniformly at random from the interval $[0, 1]$.

Healthcare MDP: The considered Healthcare MDP is borrowed from [9]. The MDP has six states corresponding to the deteriorating health condition of a patient with the last state $n = 6$ being an absorbing state representing the mortality terminal state. For each of the first five states, one can choose three inputs $m \in \{1, 2, 3\}$ corresponding to increasing levels of drug dosage for treatment. The goal is to minimize the invasiveness of the treatment while avoiding the terminal state. For the transition probabilities, we refer the reader to [9, Fig. D.1]. The cost function is chosen to be $c(n, m) = \sum_{n=1}^6 n^+ \mathbb{P}(n^+ | n, m) + m$ for each $n \in \{1, 2, 3, 4, 5\}$ and $m \in \{1, 2, 3\}$ and $c(6, 1) = 50$.

Graph MDP: The considered Graph MDP is borrowed from [21]. The MDP has 18 state-action pairs in total and corresponds to a simple path-finding problem. We refer the reader to [21, Sect. 3] for the description of the MDP.

B. Accelerated VI and QL Algorithms

The update rules are as follows.

(1) Nesterov accelerated VI [9]—with $v_{-1} = v_0$

$$y_k = v_k + \gamma^{-1}(1 - \sqrt{1 - \gamma^2})(v_k - v_{k-1})$$

$$v_{k+1} = y_k - (1 + \gamma)^{-1}(y_k - T(y_k)).$$

(2) Anderson accelerated VI [47]—with $v_{-1} = v_0$

$$y_k = v_k - v_{k-1}$$

$$z_k = T(v_k) - T(v_{k-1})$$

$$\delta_k = \frac{y_k^\top (v_k - T(v_k))}{y_k^\top (y_k - z_k)} \text{ if } y_k^\top (y_k - z_k) \neq 0, = 0 \text{ o.w.}$$

$$v_{k+1} = (1 - \delta_k)T(v_k) + \delta_k T(v_{k-1}).$$

(3) SQL [14]—with $q_{-1} = q_0$ and $\alpha_k = \frac{1}{k+1}$

$$\begin{cases} d_k = \hat{T}(q_k) - \hat{T}(q_{k-1}) \\ q_{k+1} = q_k - \alpha_k (q_k - \hat{T}(q_{k-1})) + (1 - \alpha_k)d_k. \end{cases}$$

(4) Zap QL [21]—with $D_{-1} = \mathbf{0}$ and $\alpha_k = \beta_k = \frac{1}{k+1}$

$$\begin{cases} D_k = (1 - \beta_k)D_{k-1} + \beta_k (I - \gamma \hat{P}(q_k)) \\ q_{k+1} = q_k - \alpha_k D_k^{-1} (q_k - \hat{T}(q_k)) \end{cases}$$

where $\hat{P}(q)$ is the synchronously sampled state-action transition matrix of the Markov chain under the greedy policy π_q w.r.t. q with elements $[\hat{P}(q)]((s, a), (s', a')) = 1$ if $s' = \hat{s}^+$, $a' = \pi_q(\hat{s}^+)$ and $= 0$, otherwise, for each $(s, a), (s', a') \in \mathcal{S} \times \mathcal{A}$, where $\hat{s}^+ \sim \mathbb{P}(\cdot | s, a)$ is again a *sample* of the next state drawn from the distribution $\mathbb{P}(\cdot | s, a)$ for the state-action pair (s, a) .

REFERENCES

- [1] R. Bellman, "A Markovian decision process," *J. Math. Mechan.*, vol. 6, no. 5, pp. 679–684, 1957.
- [2] M. L. Puterman and S. L. Brumelle, "On the convergence of policy iteration in stationary dynamic programming," *Math. Operations Res.*, vol. 4, no. 1, pp. 60–69, 1979.
- [3] R. A. Howard, *Dynamic Programming and Markov Processes*. Hoboken, NJ, USA: Wiley, 1960.
- [4] J. Grand-Clément, "From convex optimization to MDPs: A review of first-order, second-order and quasi-Newton methods for MDPs," 2021, *arXiv:2104.10677*.
- [5] N. Vieillard, O. Pietquin, and M. Geist, "On connections between constrained optimization and reinforcement learning," 2019, *arXiv:1910.08476*.
- [6] B. T. Polyak, "Some methods of speeding up the convergence of iteration methods," *USSR Comput. Math. Math. Phys.*, vol. 4, no. 5, pp. 1–17, 1964.
- [7] Y. E. Nesterov, "A method for solving the convex programming problem with convergence rate $o(1/k^2)$," *Doklady Akademii Nauk SSSR*, vol. 269, no. 3, pp. 543–547, 1983.
- [8] D. G. Anderson, "Iterative procedures for nonlinear integral equations," *J. ACM (JACM)*, vol. 12, no. 4, pp. 547–560, 1965.
- [9] V. Goyal and J. Grand-Clément, "A first-order approach to accelerated value iteration," *Operations Res.*, vol. 71, no. 2, pp. 517–535, 2022.
- [10] J. Zhang, B. O'Donoghue, and S. Boyd, "Globally convergent type-I Anderson acceleration for nonsmooth fixed-point iterations," *SIAM J. Optim.*, vol. 30, no. 4, pp. 3170–3197, 2020.
- [11] B. Halpern, "Fixed points of nonexpanding maps," *Bull. Amer. Math. Soc.*, vol. 73, no. 6, pp. 957–961, 1967.
- [12] J. Lee and E. Ryu, "Accelerating value iteration with anchoring," in *Proc. Adv. Neural Inf. Process. Syst.*, 2024, pp. 53924–53963.
- [13] C. Kamanchi, R. B. Diddigi, and S. Bhatnagar, "Generalized second order value iteration in Markov decision processes," *IEEE Trans. Autom. Control*, vol. 67, no. 8, pp. 4241–4247, Aug. 2022.
- [14] M. Ghavamzadeh, H. Kappen, M. Azar, and R. Munos, "Speedy Q-learning," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 24, 2011.
- [15] B. Weng, H. Xiong, L. Zhao, Y. Liang, and W. Zhang, "Finite-time theory for momentum Q-learning," in *Proc. 37th Conf. Uncertainty Artif. Intell.*, 2021, pp. 665–674.
- [16] A. M. Devraj, A. Bušić, and S. Meyn, "On matrix momentum stochastic approximation and applications to Q-learning," in *Proc. 57th Annu. Allerton Conf. Commun., Control, Comput.*, 2019, pp. 749–756.
- [17] T. Yang, Q. Lin, and Z. Li, "Unified convergence analysis of stochastic momentum methods for convex and non-convex optimization," 2016, *arXiv:1604.03257*.
- [18] R. Kidambi, P. Netrapalli, P. Jain, and S. Kakade, "On the insufficiency of existing momentum schemes for stochastic optimization," in *Inf. Theory Appl. Workshop*, 2018, pp. 1–9.
- [19] C. Liu and M. Belkin, "Accelerating SGD with momentum for over-parameterized learning," 2018, *arXiv:1810.13395*.
- [20] Z. Allen-Zhu, "Katyusha: The first direct acceleration of stochastic gradient methods," *J. Mach. Learn. Res.*, vol. 18, no. 1, pp. 8194–8244, 2017.

- [21] A. M. Devraj and S. Meyn, "Zap Q-learning," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 30, 2017.
- [22] D. Ruppert, "A Newton-Raphson version of the multivariate Robbins-Monro procedure," *Ann. Statist.*, vol. 13, no. 1, pp. 236–245, 1985.
- [23] D. Bertsekas, "Convergence of discretization procedures in dynamic programming," *IEEE Trans. Autom. Control*, vol. 20, no. 3, pp. 415–419, Jun. 1975.
- [24] W. B. Powell, *Approximate Dynamic Programming: Solving the Curses of Dimensionality*. Hoboken, NJ, USA: Wiley, 2007.
- [25] D. Bertsekas, *Abstract Dynamic Programming*. Belmont, MA, USA: Athena Sci., 2022.
- [26] C. Szepesvári, *Algorithms for Reinforcement Learning*. San Rafael, CA, USA: Morgan & Claypool, 2010.
- [27] D. P. Bertsekas, "Temporal difference methods for general projected equations," *IEEE Trans. Autom. Control*, vol. 56, no. 9, pp. 2128–2139, Sep. 2011.
- [28] J. Tsitsiklis and B. Van Roy, "An analysis of temporal-difference learning with function approximation," *IEEE Trans. Autom. Control*, vol. 42, no. 5, pp. 674–690, May 1997.
- [29] D. Bertsekas and J. N. Tsitsiklis, *Neuro-Dynamic Program.* Belmont, MA, USA: Athena Sci., 1996.
- [30] J. Schmidhuber, "Deep learning in neural networks: An overview," *Neural Netw.*, vol. 61, pp. 85–117, 2015.
- [31] R. S. Sutton and A. G. Barto, *Reinforcement Learning: An Introduction*. Cambridge, MA, USA: MIT Press, 2018.
- [32] E. Berthier and F. Bach, "Max-plus linear approximations for deterministic continuous-state Markov decision processes," *IEEE Control Syst. Lett.*, vol. 4, no. 3, pp. 767–772, Jul. 2020.
- [33] V. M. Gonçalves, "Max-plus approximation for reinforcement learning," *Automatica*, vol. 129, 2021, Art. no. 109623.
- [34] M. A. S. Kolarijani and P. Mohajerin Esfahani, "Fast approximate dynamic programming for input-affine dynamics," *IEEE Trans. Autom. Control*, vol. 68, no. 10, pp. 6315–6322, Oct. 2023.
- [35] M. A. S. Kolarijani, G. F. Max, and P. Mohajerin Esfahani, "Fast approximate dynamic programming for infinite-horizon Markov decision processes," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 34, 2021, pp. 23652–23663.
- [36] W. M. McEneaney, *Max-Plus Methods for Nonlinear Control and Estimation*. Berlin, Germany: Springer, 2006.
- [37] Y. Liu and M. A. S. Kolarijani, "Fitted Q-iteration via max-plus-linear approximation," *IEEE Control Syst. Lett.*, vol. 8, pp. 3201–3206, 2024.
- [38] O. Hernández-Lerma and J. B. Lasserre, *Discrete-Time Markov Control Processes: Basic Optimality Criteria*. Berlin, Germany: Springer, 2012.
- [39] D. P. De Farias and B. Van Roy, "On constraint sampling in the linear programming approach to approximate dynamic programming," *Math. Operations Res.*, vol. 29, no. 3, pp. 462–478, 2004.
- [40] O. Hernández-Lerma and J. B. Lasserre, *Further Topics on Discrete-Time Markov Control Processes*. Berlin, Germany: Springer, 2012.
- [41] P. Mohajerin Esfahani, T. Sutter, D. Kuhn, and J. Lygeros, "From infinite to finite programs: Explicit error bounds with applications to approximate dynamic programming," *SIAM J. Optim.*, vol. 28, no. 3, pp. 1968–1998, 2018.
- [42] M. L. Puterman, *Markov Decision Processes: Discrete Stochastic Dynamic Programming*. Hoboken, NJ, USA: Wiley, 2014.
- [43] D. Bertsekas, *Lessons from AlphaZero for Optimal, Model Predictive, and Adaptive Control*. Belmont, MA, USA: Athena Sci., 2022.
- [44] M. Gargiani, A. Zanelli, D. Liao-McPherson, T. Summers, and J. Lygeros, "Dynamic programming through the lens of semismooth Newton-type methods," *IEEE Control Syst. Lett.*, vol. 6, pp. 2996–3001, 2022.
- [45] S. Bubeck, "Convex optimization: Algorithms and complexity," *Found. Trends Mach. Learn.*, vol. 8, no. 3–4, pp. 231–357, 2015.
- [46] C. G. Broyden, "A class of methods for solving nonlinear simultaneous equations," *Math. Computation*, vol. 19, no. 92, pp. 577–593, 1965.
- [47] M. Geist and B. Scherrer, "Anderson acceleration for reinforcement learning," 2018, *arXiv:1809.09501*.
- [48] K. Sun et al., "Damped anderson mixing for deep reinforcement learning: Acceleration, convergence, and stabilization," in *Proc. Adv. Neural Inf. Process. Syst.*, 2021, vol. 34, pp. 3732–3743.
- [49] M. Kearns and S. Singh, "Finite-sample convergence rates for Q-learning and indirect algorithms," in *Proc. Adv. Neural Inf. Process. Syst.*, 1998, pp. 996–1002.
- [50] C. J. Watkins and P. Dayan, "Q-learning," *Mach. Learn.*, vol. 8, no. 3, pp. 279–292, 1992.
- [51] F. Facchinei and J. Pang, *Finite-Dimensional Variational Inequalities and Complementarity Problems*. Berlin, Germany: Springer, 2003, vol. 2.
- [52] D. Sun and J. Han, "Newton and quasi-Newton methods for a class of nonsmooth equations and related problems," *SIAM J. Optim.*, vol. 7, no. 2, pp. 463–480, 1997.
- [53] S. Perkins and D. S. Leslie, "Asynchronous stochastic approximation with differential inclusions," *Stochastic Syst.*, vol. 2, no. 2, pp. 409–446, 2013.
- [54] T. Archibald, K. McKinnon, and L. Thomas, "On the generation of Markov decision processes," *J. Oper. Res. Soc.*, vol. 46, no. 3, pp. 354–361, 1995.
- [55] A. Rakhsha, A. Wang, M. Ghavamzadeh, and A.-M. Farahmand, "Operator splitting value iteration," in *Proc. Adv. Neural Inf. Process. Syst.*, 2022, vol. 35, pp. 38373–38385.
- [56] C. Evans, S. Pollock, L. G. Rebholz, and M. Xiao, "A proof that anderson acceleration improves the convergence rate in linearly converging fixed-point methods (but not in those converging quadratically)," *SIAM J. Numer. Anal.*, vol. 58, no. 1, pp. 788–810, 2020.



Mohamad Amin Sharifi Kolarijani received the B.Sc. degree in control systems from the University of Tehran, Tehran, Iran, in 2011, the M.Sc. degree in bio-electrics from the Amirkabir University of Technology, Tehran, Iran, in 2014, and the Ph.D. degree in systems and control from the Delft University of Technology, Delft, The Netherlands, in 2022.

He is currently an Assistant Professor with the Delft Center for Systems and Control, Delft University of Technology. His research is mainly

focused on developing algorithms for solving dynamic programming and reinforcement learning problems using tools from computational convex analysis, max-plus algebra, and algorithmic optimization.

Dr. Kolarijani was the recipient of the Bronze Medal of the Iranian National Olympiad in Physics in 2005, and was ranked 3rd in Iran's National Matriculation Exam for Graduate Studies in 2017.



Peyman Mohajerin Esfahani received the B.Sc. and M.Sc. degrees in electrical engineering from the Sharif University of Technology, Tehran, Iran, in 2005 and 2008, and the Ph.D. degree in automatic control from ETH Zurich, Zurich, Switzerland, in 2014.

He is currently an Associate Professor with the University of Toronto, Toronto, ON, Canada, and the Delft University of Technology (TU Delft), Delft, The Netherlands. He joined TU Delft in October 2016, and prior to that, he held

several research appointments at EPFL, ETH Zurich, and MIT between 2014 and 2016. His research interests include theoretical and practical aspects of decision-making problems in uncertain and dynamic environments, with applications in engineering systems and management.

Dr. Esfahani serves as an Associate Editor for IEEE TRANSACTIONS ON AUTOMATIC CONTROL, *Operations Research*, *Mathematical Programming*, and *Open Journal of Mathematical Optimization*. He was the recipient of the ERC Starting Grant, one of the three finalists for the Young Researcher Prize in Continuous Optimization awarded by the Mathematical Optimization Society in 2016 and a recipient of the 2016 George S. Axelby Outstanding Paper Award from the IEEE Control Systems Society, the INFORMS Frederick W. Lanchester Prize in 2020, and the 2022 European Control Award.