

## Technical Appendix: Proofs, Algorithms, and Supplementary Results

### EC.1. Summary of Notation

|                                                |                                                                                                                                                                                                                                                                                                  |
|------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $A$                                            | consumption matrix, $A := [a_{ij}]_{m \times n}$                                                                                                                                                                                                                                                 |
| $A^j$                                          | $j$ -th column of $A$ , incidence vector for product $j$                                                                                                                                                                                                                                         |
| $\mathcal{A}$                                  | action space, collection of all subsets of $\mathcal{J}$                                                                                                                                                                                                                                         |
| $\mathcal{A}(x)$                               | collection of all available assortments at the state $x$ ,<br>$\mathcal{A}(x) := \{S \in \mathcal{A} : x \geq A^j \text{ for all } j \in S\}$                                                                                                                                                    |
| $B([0, T] \times \mathcal{X})$                 | space of Borel measurable functions defined on $[0, T] \times \mathcal{X}$ such that<br>$\ \psi\  := \sup_{(t,x) \in [0,T] \times \mathcal{X}}  \psi(t,x)  < \infty$                                                                                                                             |
| $c$                                            | initial inventory of $m$ resources, $c = (c_1, \dots, c_m)^\top$                                                                                                                                                                                                                                 |
| $C^{1,0}([0, T] \times \mathcal{X})$           | space of real-valued functions defined on $[0, T] \times \mathcal{X}$ that are continuously differentiable in $t$ over $[0, T]$ for all $x \in \mathcal{X}$                                                                                                                                      |
| $d$                                            | degree of polynomial involved in parametric families of the value functions or policies                                                                                                                                                                                                          |
| $\mathcal{F}$                                  | rich enough $\sigma$ -algebra defined on the sample space $\Omega$                                                                                                                                                                                                                               |
| $\{\mathcal{F}_t\}_{t \geq 0}$                 | filtration generated by standard Poisson process $N^\lambda$ and two sequences of i.i.d. uniform random variables for randomized actions and customer choices, respectively                                                                                                                      |
| $\{\mathcal{F}_t^{X^\pi}\}_{t \geq 0}$         | natural filtration generated by state process $X^\pi$ under randomized policy $\pi$                                                                                                                                                                                                              |
| $\{\mathcal{F}_t^{\tilde{X}^\pi}\}_{t \geq 0}$ | natural filtration generated by exploratory state process $\tilde{X}^\pi$                                                                                                                                                                                                                        |
| $\{\mathcal{G}_t\}_{t \geq 0}$                 | filtration generated by standard Poisson process $N^\lambda$ and a sequences of i.i.d. uniform random variables for customer choices                                                                                                                                                             |
| $H(t, x, S, v(\cdot, \cdot))$                  | Hamiltonian of the classical intensity control problem                                                                                                                                                                                                                                           |
| $\mathcal{H}(\pi(\cdot   t, x))$               | entropy of randomized policy $\pi(\cdot   t, x)$                                                                                                                                                                                                                                                 |
| $J(t, x; \pi)$                                 | entropy-regularized value function under policy $\pi$ in continuous-time RL formulation                                                                                                                                                                                                          |
| $J_{\Delta t}(t_k, x; \pi)$                    | entropy-regularized value function under policy $\pi$ in discretization-based RL methods                                                                                                                                                                                                         |
| $J^*(t, x)$                                    | optimal value function in the continuous-time RL formulation                                                                                                                                                                                                                                     |
| $J^\theta(t, x)$                               | parametric value function in continuous-time RL formulation                                                                                                                                                                                                                                      |
| $J_{\Delta t}^\theta(t_k, x)$                  | parametric value function in discretization-based RL methods                                                                                                                                                                                                                                     |
| $\mathcal{J}$                                  | set of all products, $\mathcal{J} = \{1, \dots, n\}$                                                                                                                                                                                                                                             |
| $K$                                            | total number of intervals in a uniform grid $0 = t_0 < t_1 < \dots < t_K = T$ with step size $\Delta t = \frac{T}{K}$                                                                                                                                                                            |
| $\mathbb{K}$                                   | set of all valid time-state-action triples, $\mathbb{K} := \{(t, x, S) : t \in [0, T], x \in \mathcal{X}, S \in \mathcal{A}(x)\}$                                                                                                                                                                |
| $L$                                            | total number of state jumps over $[0, T]$                                                                                                                                                                                                                                                        |
| $\bar{L}$                                      | number of disjoint consideration sets for customers                                                                                                                                                                                                                                              |
| $m$                                            | number of resources                                                                                                                                                                                                                                                                              |
| $M$                                            | batch size in algorithms                                                                                                                                                                                                                                                                         |
| $M_{\Delta t}$                                 | discrete-time model under a uniform grid $0 = t_0 < t_1 < \dots < t_K = T$ with step size $\Delta t = \frac{T}{K}$ , it is assumed that in each period $(t_k, t_{k+1}]$ , one arrival occurs with probability $\lambda \Delta t$ , and no arrivals occur with probability $1 - \lambda \Delta t$ |
| $n$                                            | number of products                                                                                                                                                                                                                                                                               |
| $\bar{n}$                                      | total number of revenue-ordered assortments                                                                                                                                                                                                                                                      |
| $N$                                            | number of episodes in algorithms                                                                                                                                                                                                                                                                 |
| $N^\lambda, N_t^\lambda$                       | a standard Poisson process with rate $\lambda$ , to model the arrival process of all potential customers                                                                                                                                                                                         |
| $N^u, N_t^u$                                   | $N_t^u = (N_{1,t}^u, \dots, N_{n,t}^u)^\top$ , a vector of Poisson processes with intensities $(\lambda P_1(S_t), \dots, \lambda P_n(S_t))$ modulated by deterministic control $u = \{S_t : t \in [0, T]\}$                                                                                      |

|                                                      |                                                                                                                                                                                                                                                                                                   |
|------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| $N^\pi, N_t^\pi$                                     | $N_t^\pi = (N_{1,t}^\pi, \dots, N_{n,t}^\pi)^\top$ , with $N_{j,t}^\pi$ representing the number of product $j$ sold up to time $t$ under the randomized policy $\pi$                                                                                                                              |
| $p$                                                  | prices vector for $n$ products, $p = (p_1, \dots, p_n)^\top$                                                                                                                                                                                                                                      |
| $P_0(S)$                                             | probability that an arriving customer makes no purchase when assortment $S$ is offered                                                                                                                                                                                                            |
| $P_j(S)$                                             | probability that an arriving customer purchases product $j$ when assortment $S$ is offered                                                                                                                                                                                                        |
| $\mathbb{P}$                                         | probability measure defined on $\mathcal{F}$                                                                                                                                                                                                                                                      |
| $q(y   t, x, S)$                                     | transition rate of state process                                                                                                                                                                                                                                                                  |
| $r(S)$                                               | average revenue rate when offered $S$ , $r(S) := \lambda \sum_{j=1}^n p_j P_j(S)$                                                                                                                                                                                                                 |
| $S$                                                  | assortment offered by the firm, $S \subseteq \mathcal{J}$                                                                                                                                                                                                                                         |
| $S_t^\pi$                                            | assortment offered at time $t$ under policy $\pi$                                                                                                                                                                                                                                                 |
| $T$                                                  | selling time limit                                                                                                                                                                                                                                                                                |
| $\mathbf{u}$                                         | deterministic control process, $\mathbf{u} = \{S_t \in \mathcal{A} : t \in [0, T]\}$                                                                                                                                                                                                              |
| $\mathcal{U}$                                        | set of all non-anticipating control processes                                                                                                                                                                                                                                                     |
| $V(t, x; \mathbf{u})$                                | value function under control $\mathbf{u}$ in the classical formulation                                                                                                                                                                                                                            |
| $V^*(t, x)$                                          | optimal value function in the classical formulation                                                                                                                                                                                                                                               |
| $V_{\Delta t}^*(t_k, x)$                             | optimal value function under the discrete-time model $\mathcal{M}_{\Delta t}$                                                                                                                                                                                                                     |
| $V_{\text{ADP-}\Delta t}(t_k, x)$                    | approximate value function derived from ADP under the discrete-time model $\mathcal{M}_{\Delta t}$ , $V_{\text{ADP-}\Delta t}(t_k, x) = \theta_k + \sum_{i=1}^m V_{k,i} x_i$                                                                                                                      |
| $V_{\text{I-ADP-}\Delta t}(t, x)$                    | continuous-time interpolated approximated value function based on $V_{\text{ADP-}\Delta t}(t_k, x)$                                                                                                                                                                                               |
| $X^u, X_t^u, X_{t-}^u$                               | process of the remaining inventory (state) under deterministic control $\mathbf{u}$ , $X_{t-}^u$ refers to the left limit of process $X^u$ at time $t$                                                                                                                                            |
| $X^\pi, X_t^\pi, X_{t-}^\pi$                         | process of the remaining inventory (state) under randomized policy $\pi$ , $X_{t-}^\pi$ refers to the left limit of process $X^\pi$ at time $t$                                                                                                                                                   |
| $\tilde{X}^\pi, \tilde{X}_t^\pi, \tilde{X}_{t-}^\pi$ | exploratory state process defined on the probability space $(\Omega, \mathcal{F}, \mathbb{P}; \{\mathcal{G}_t\}_{t \geq 0})$ , it has the same distribution as the sample state process $\{X_t^\pi : t \in [0, T]\}$ defined on $(\Omega, \mathcal{F}, \mathbb{P}; \{\mathcal{F}_t\}_{t \geq 0})$ |
| $X$                                                  | state space, $X = \{0, \dots, c_1\} \times \dots \times \{0, \dots, c_m\}$                                                                                                                                                                                                                        |
| $\alpha_\phi$                                        | learning rate for actor (policy) network                                                                                                                                                                                                                                                          |
| $\alpha_\theta$                                      | learning rate for critic (value function) network                                                                                                                                                                                                                                                 |
| $\gamma$                                             | entropy factor (temperature parameter)                                                                                                                                                                                                                                                            |
| $\lambda$                                            | customer arrival rate                                                                                                                                                                                                                                                                             |
| $\pi(S   t, x)$                                      | randomized policy                                                                                                                                                                                                                                                                                 |
| $\pi^*(S   t, x)$                                    | optimal randomized policy for the entropy-regularized optimization problem                                                                                                                                                                                                                        |
| $\pi^\phi(S   t, x)$                                 | parametric randomized policy                                                                                                                                                                                                                                                                      |
| $\sigma^\phi(S^{[j]}   t)$                           | parametric policy with action space that comprises all revenue-ordered assortments $S^{[1]}, \dots, S^{[\bar{n}]}$                                                                                                                                                                                |
| $\varphi(t, x)$                                      | vector of basis functions for the linear value function approximation, $\varphi(t, x) := (\varphi_1(t, x), \dots, \varphi_W(t, x))^\top$                                                                                                                                                          |
| $\phi$                                               | parameters in actor (policy) network                                                                                                                                                                                                                                                              |
| $\theta$                                             | parameters in critic (value function) network                                                                                                                                                                                                                                                     |
| $\tau_l$                                             | $l$ -th jump time the state process                                                                                                                                                                                                                                                               |

## EC.2. The Reformulation of (6) in Section 2.2

We first prove that for each  $j = 1, \dots, n$ , the process  $\{N_{j,t}^\pi - \int_0^t \sum_{S \in \mathcal{A}(X_{t-}^\pi)} \lambda P_j(S) \pi(S | s, X_{t-}^\pi) dt : t \in [0, T]\}$  is an  $\{\mathcal{F}_t\}_{t \geq 0}$ -martingale. Let  $N_{(t, t+\Delta t]}^\lambda$  denote the number of customer arrivals in  $(t, t + \Delta t]$ ; if there is exactly one customer arrival during this time interval, we denote the arrival time by  $t_a$ . Given  $N_{(t, t+\Delta t]}^\lambda = 1$ , the random variable  $t_a$  is uniformly distributed over  $(t, t + \Delta t]$  by the property of the Poisson process  $N^\lambda$ . To

demonstrate the Markov property of  $N^\pi$  and derive its infinitesimal generator, for any bounded, measurable, real-valued function  $f$ , we evaluate  $\mathbb{E}[f(N_{t+\Delta t}^\pi) \mid N_t^\pi = z, \mathcal{F}_t]$  as follows:

$$\begin{aligned}
& \mathbb{E}[f(N_{t+\Delta t}^\pi) \mid N_t^\pi = z, \mathcal{F}_t] \\
&= \sum_{y \in \mathcal{X}} f(y) \cdot \mathbb{P}(N_{t+\Delta t}^\pi = y \mid N_t^\pi = z, \mathcal{F}_t) \\
&= \sum_{y \in \mathcal{X}} f(y) \cdot \left\{ \mathbb{P}(N_{t+\Delta t}^\pi = y \mid N_{(t,t+\Delta t]}^\lambda = 0, N_t^\pi = z, \mathcal{F}_t) \cdot \mathbb{P}(N_{(t,t+\Delta t]}^\lambda = 0 \mid N_t^\pi = z, \mathcal{F}_t) \right. \\
&\quad + \left[ \int_t^{t+\Delta t} \mathbb{P}(N_{t+\Delta t}^\pi = y \mid t_a = u, N_{(t,t+\Delta t]}^\lambda = 1, N_t^\pi = z, \mathcal{F}_t) \cdot \frac{1}{\Delta t} du \right] \cdot \mathbb{P}(N_{(t,t+\Delta t]}^\lambda = 1 \mid N_t^\pi = z, \mathcal{F}_t) \\
&\quad \left. + \mathbb{P}(N_{t+\Delta t}^\pi = y \mid N_{(t,t+\Delta t]}^\lambda \geq 2, N_t^\pi = z, \mathcal{F}_t) \cdot \mathbb{P}(N_{(t,t+\Delta t]}^\lambda \geq 2 \mid N_t^\pi = z, \mathcal{F}_t) \right\} \\
&= f(z) \cdot [1 - \lambda \Delta t + o(\Delta t)] + \left[ f(z) \left( \sum_{S \in \mathcal{A}(c-Az)} P_0(S) \frac{\int_t^{t+\Delta t} \pi(S \mid u, c-Az) du}{\Delta t} \right) \right. \\
&\quad \left. + \sum_{j=1}^n f(z + e^j) \left( \sum_{S \in \mathcal{A}(c-Az)} P_j(S) \frac{\int_t^{t+\Delta t} \pi(S \mid u, c-Az) du}{\Delta t} \right) \right] \cdot [\lambda \Delta t + o(\Delta t)] + o(\Delta t) \\
&= f(z) + \lambda \Delta t \cdot \sum_{j=1}^n \left( \sum_{S \in \mathcal{A}(c-Az)} P_j(S) \frac{\int_t^{t+\Delta t} \pi(S \mid u, c-Az) du}{\Delta t} \right) \cdot [f(z + e^j) - f(z)] + o(\Delta t).
\end{aligned}$$

Since the result depends only on the current state  $z$ , the process  $N^\pi$  satisfies the Markov property. Next, its infinitesimal generator is given by

$$\lim_{\Delta t \rightarrow 0} \frac{\mathbb{E}[f(N_{t+\Delta t}^\pi) \mid N_t^\pi = z] - f(z)}{\Delta t} = \lambda \sum_{j=1}^n \sum_{S \in \mathcal{A}(c-Az)} P_j(S) \pi(S \mid t, c-Az) \cdot [f(z + e^j) - f(z)]. \quad (\text{EC.1})$$

Given that  $N^\pi$  has generator (EC.1) and  $X_t^\pi = c - AN_t^\pi$  is bounded for all  $t \in [0, T]$ , it follows from Dynkin's formula that for each  $j = 1, \dots, n$ , the process  $\{N_{j,t}^\pi - \int_0^t \sum_{S \in \mathcal{A}(X_{t-}^\pi)} \lambda P_j(S) \pi(S \mid s, X_{t-}^\pi) dt : t \in [0, T]\}$  is an  $\{\mathcal{F}_t\}_{t \geq 0}$ -martingale.

Next, we introduce the exploratory dynamics of  $\{\tilde{X}_t^\pi : t \in [0, T]\}$ , which can be interpreted as the average of trajectories of the sample state process  $\{X_t^\pi : t \in [0, T]\}$  over infinitely many randomized actions. In the case of controlled diffusions, Wang et al. (2020) derived heuristically the exploratory state process by applying a law-of-large-numbers argument to the drift and diffusion coefficients of the controlled diffusion process. However, the sample state process  $\{X_t^\pi : t \in [0, T]\}$  in our setting is a pure jump process, so their approach does not apply. Instead, we derive the infinitesimal generator of the sample state process  $\{X_t^\pi : t \in [0, T]\}$ , from which we will identify the dynamics of the exploratory state process.

From the above derivation for the generator of  $N^\pi$ , it is evident that for  $s < t$ , given the present state  $X_s^\pi$ , the increment  $X_t^\pi - X_s^\pi = -A(N_t^\pi - N_s^\pi)$  does not depend on any other past information in  $\mathcal{F}_s$ . Therefore, the sample state process  $\{X_t^\pi : t \in [0, T]\}$  is also a Markov process. Using a similar derivation as for (EC.1), we

can obtain the generator of  $\{X_t^\pi : t \in [0, T]\}$ . For each time-state pair  $(t, x)$  and any bounded and measurable function  $g : \mathcal{X} \mapsto \mathbb{R}$ ,

$$\lim_{\Delta t \rightarrow 0} \frac{\mathbb{E}[g(X_{t+\Delta t}^\pi) | X_t^\pi = x] - g(x)}{\Delta t} = \lambda \sum_{j=1}^n \sum_{S \in \mathcal{A}(x)} P_j(S) \pi(S | t, x) \cdot [g(x - A^j) - g(x)]. \quad (\text{EC.2})$$

One can observe from (EC.2) that the effect of individually sampled actions has been averaged out. We can construct an (“averaged”) exploratory process  $\{\tilde{X}_t^\pi : t \in [0, T]\}$ , defined on the original probability space  $(\Omega, \mathcal{F}, \mathbb{P}; \{\mathcal{G}_t\}_{t \geq 0})$ , as a continuous-time Markov chain with the generator given by (EC.2) and  $\tilde{X}_0^\pi = c$ . Because the generator and the initial state of  $\{X_t^\pi : t \in [0, T]\}$  and  $\{\tilde{X}_t^\pi : t \in [0, T]\}$  are identical, we infer that the sample state process  $\{X_t^\pi : t \in [0, T]\}$  has the same distribution as the distribution of the exploratory state process  $\{\tilde{X}_t^\pi : t \in [0, T]\}$ . This completes the proof of (6).

### EC.3. Actor-Critic Algorithms

### EC.4. Proofs of Statements

*Proof of Lemma 1* Let  $B([0, T] \times \mathcal{X})$  denote the space of all Borel measurable functions  $\psi : [0, T] \times \mathcal{X} \rightarrow \mathbb{R}$  such that  $\|\psi\| := \sup_{(t,x) \in [0,T] \times \mathcal{X}} |\psi(t, x)| < \infty$ . This space becomes a Banach space when endowed with the supreme norm  $\|\cdot\|$ . We define an operator  $\mathcal{L}$  on  $B([0, T] \times \mathcal{X})$  by

$$\mathcal{L}\psi(t, x) := e^{\beta t} \int_t^T \left\{ R^\pi(s, x) + \sum_{S \in \mathcal{A}(x)} \left( \sum_{y \in \mathcal{X}} e^{-\beta s} \psi(s, y) q(y | s, x, S) \right) \pi(S | s, x) \right\} ds, \quad (\text{EC.3})$$

where  $R^\pi(t, x)$  is as defined in (5), and  $\beta := 2\lambda + 1$ . Given that the integrand on the right-hand side of (EC.3) also belongs to the space  $B([0, T] \times \mathcal{X})$ , it follows that the operator  $\mathcal{L}$  is well-defined. Then, for  $\psi_1, \psi_2 \in B([0, T] \times \mathcal{X})$ , we obtain

$$\begin{aligned} |\mathcal{L}\psi_1(t, x) - \mathcal{L}\psi_2(t, x)| &\leq e^{\beta t} \int_t^T e^{-\beta s} \left( \sum_{S \in \mathcal{A}(x)} \sum_{y \in \mathcal{X}} |\psi_1(s, y) - \psi_2(s, y)| \cdot |q(y | s, x, S)| \pi(S | s, x) \right) ds \\ &\leq e^{\beta t} \int_t^T e^{-\beta s} \left( \|\psi_1 - \psi_2\| \cdot \sum_{S \in \mathcal{A}(x)} \left( \sum_{y \in \mathcal{X}} |q(y | s, x, S)| \right) \pi(S | s, x) \right) ds \\ &\leq 2\lambda \|\psi_1 - \psi_2\| e^{\beta t} \int_t^T e^{-\beta s} ds \\ &\leq \frac{2\lambda}{\beta} (1 - e^{-\beta(T-t)}) \|\psi_1 - \psi_2\| \\ &\leq \frac{2\lambda}{\beta} \|\psi_1 - \psi_2\|. \end{aligned}$$

Observing that  $\frac{2\lambda}{\beta} = \frac{2\lambda}{2\lambda+1} < 1$ , we identify  $\mathcal{L}$  as a contraction operator on the Banach space  $B([0, T] \times \mathcal{X})$ .

Let  $\psi^* \in B([0, T] \times \mathcal{X})$  be the fixed point of  $\mathcal{L}$ , i.e.,

$$\psi^*(t, x) := e^{\beta t} \int_t^T \left\{ R^\pi(s, x) + \sum_{S \in \mathcal{A}(x)} \left( \sum_{y \in \mathcal{X}} e^{-\beta s} \psi^*(s, y) q(y | s, x, S) \right) \pi(S | s, x) \right\} ds.$$

**Algorithm EC.1** Actor-Critic Algorithm with Linear Value Function Approximation (PE via TD(0))

- 
- 1: **Inputs:** initial state  $c$ , time horizon  $T$ , number of episodes  $N$ , batch size  $M$ ; functional forms of basis functions  $\varphi_1, \dots, \varphi_W$ , functional form of the policy  $\pi^\phi(\cdot | \cdot, \cdot)$  and an initial value  $\phi_0$ ; entropy factor  $\gamma$ , learning rate  $\alpha_\phi$
- 2: **Required program:** an environment simulator  $(t', x', S', r') = \text{Environment}(t, x, \pi^\phi(\cdot | \cdot, \cdot))$
- 3: Initialize  $\phi \leftarrow \phi_0$
- 4: **for**  $i = 1$  to  $N$  **do**
- 5:     Store  $(\tau_0^{(i)}, x_0^{(i)}) \leftarrow (0, c)$
- 6:     Initialize  $l = 0$ ,  $(t, x) = (0, c)$   $\triangleright$  Initialize  $l$  to count jumps in each state trajectory, and  $(t, x)$  to record the time and state right after each jump
- 7:     **while** True **do**
- 8:         Apply  $(t, x)$  to the environment simulator to get  $(t', x', S', r') = \text{Environment}(t, x, \pi^\phi(\cdot | \cdot, \cdot))$
- 9:         **if**  $t' \geq T$  **then**
- 10:             Store  $L^{(i)} \leftarrow l$ ,  $\tau_{L^{(i)}+1}^{(i)} \leftarrow T$
- 11:             **break**
- 12:         **end if**
- 13:         Update  $l \leftarrow l + 1$
- 14:         Store observation at jump time:  $(\tau_l^{(i)}, x_l^{(i)}, S_l^{(i)}, r_l^{(i)}) \leftarrow (t', x', S', r')$
- 15:         Update  $(t, x) \leftarrow (t', x')$
- 16:     **end while**
- 17:     Denote  $G_l^\varphi(i) := \varphi(\tau_l^{(i)}, x_{l-1}^{(i)})[\varphi(\tau_l^{(i)}, x_l^{(i)}) - \varphi(\tau_l^{(i)}, x_{l-1}^{(i)})]$ ,  $F_l^\varphi(i) := F^\varphi(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)})$
- 18:      $E_l^\varphi(\pi^\phi, i) := E(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \varphi(\cdot, x_l^{(i)}), \pi^\phi)$ ,  $E_l^1(\pi^\phi, i) = E(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \mathbf{1}, \pi^\phi)$
- 19:      $U_l^{\theta^*}(\pi^\phi, i) := \log \pi^\phi(S_l^{(i)} | \tau_l^{(i)}, x_{l-1}^{(i)})[J^{\theta^*}(\tau_l^{(i)}, x_l^{(i)}) - J^{\theta^*}(\tau_l^{(i)}, x_{l-1}^{(i)}) + r_l^{(i)}]$
- 20:     **if**  $i = 0 \pmod{M}$  **then**  $\triangleright$  Perform an update using  $M$  episodes generated under the policy
- 21:         Evaluate policy  $\pi^\phi$ : [using formula (19), incorporating techniques (20) and (21)]
- $$\theta^* = \left[ \frac{1}{M} \sum_{k=i-M+1}^i \left( \sum_{l=1}^{L^{(k)}} G_l^\varphi(k) + \sum_{l=0}^{L^{(k)}} F_l^\varphi(k) \right) \right]^{(-1)} \times$$
- $$\left[ \frac{1}{M} \sum_{k=i-M+1}^i \left( \sum_{l=1}^{L^{(k)}} \varphi(\tau_l^{(k)}, x_{l-1}^{(k)}) r_{\tau_l^{(k)}} + \gamma \sum_{l=0}^{L^{(k)}} E_l^\varphi(\pi^\phi, k) \right) \right]$$
- 22:         Compute policy gradient at  $\phi$ : [using formula (25)]
- $$\Delta\phi = \nabla_\phi \left\{ \frac{1}{M} \sum_{k=i-M+1}^i \left( \sum_{l=1}^{L^{(k)}} U_l^{\theta^*}(\pi^\phi, k) + \gamma \sum_{l=0}^{L^{(k)}} E_l^1(\pi^\phi, k) \right) \right\}$$
- 23:         Update  $\phi$  by
- $$\phi \leftarrow \phi + \alpha_\phi \Delta\phi$$
- 24:     **end if**
- 25: **end for**
-

**Algorithm EC.2** Actor-Critic Algorithm with Neural Networks (PE via Monte Carlo)

---

1: **Inputs:** initial state  $c$ , time horizon  $T$ , number of episodes  $N$ , batch size  $M$ ; critic network  $J^\theta(t, x)$  and an initial value  $\theta_0$ ; actor network  $\pi^\phi(S | t, x)$  and an initial value  $\phi_0$ ; entropy factor  $\gamma$ , learning rates  $\alpha_\theta, \alpha_\phi$

2: **Required program:** an environment simulator  $(t', x', S', r') = \text{Environment}(t, x, \pi^\phi(\cdot | \cdot, \cdot))$

3: Initialize  $\theta \leftarrow \theta_0, \phi \leftarrow \phi_0$

4: **for**  $i = 1$  to  $N$  **do**

5:     Store  $(\tau_0^{(i)}, x_0^{(i)}) \leftarrow (0, c)$

6:     Initialize  $l = 0, (t, x) = (0, c)$   $\triangleright$  Initialize  $l$  to count jumps in each state trajectory, and  $(t, x)$  to record the time and state right after each jump

7:     **while** True **do**

8:         Apply  $(t, x)$  to the environment simulator to get  $(t', x', S', r') = \text{Environment}(t, x, \pi^\phi(\cdot | \cdot, \cdot))$

9:         **if**  $t' \geq T$  **then**

10:             Store  $L^{(i)} \leftarrow l, \tau_{L^{(i)}+1}^{(i)} \leftarrow T$

11:             **break**

12:         **end if**

13:         Update  $l \leftarrow l + 1$

14:         Store observation at jump time:  $(\tau_l^{(i)}, x_l^{(i)}, S_l^{(i)}, r_l^{(i)}) \leftarrow (t', x', S', r')$

15:         Update  $(t, x) \leftarrow (t', x')$

16:     **end while**

17:     Denote  $D_l(\theta, i) := D(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \theta), b_l(\theta, i) := b(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \theta)$

18:      $C_l(\pi^\phi, i) := \sum_{l'=l+1}^{L^{(i)}} [r_{l'}^{(i)} + \gamma E(\tau_{l'}^{(i)}, \tau_{l'+1}^{(i)}, x_{l'}^{(i)}; \mathbf{1}, \pi^\phi)]$

19:      $E_l(\theta, \pi^\phi, i) := E(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; b(\tau_l^{(i)}, \cdot, x_l^{(i)}; \theta), \pi^\phi), E_l^1(\pi^\phi, i) = E(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \mathbf{1}, \pi^\phi)$

20:      $U_l^\theta(\pi^\phi, i) := \log \pi^\phi(S_l^{(i)} | \tau_l^{(i)}, x_{l-1}^{(i)}) [J^\theta(\tau_l^{(i)}, x_l^{(i)}) - J^\theta(\tau_l^{(i)}, x_{l-1}^{(i)}) + r_l^{(i)}]$

21:

22:     **if**  $i = 0 \pmod{M}$  **then**  $\triangleright$  Perform an update using  $M$  episodes generated under the policy

23:         Update critic network: [using formula (14)]

$$\theta \leftarrow \theta - \alpha_\theta \nabla_\theta \left\{ \frac{1}{M} \sum_{k=i-M+1}^i \sum_{l=0}^{L^{(k)}} \left( \frac{1}{2} D_l(\theta, k) - b_l(\theta, k) C_l(\pi^\phi, k) - \gamma E_l(\theta, \pi^\phi, k) \right) \right\}$$

24:         Update actor network: [using formula (25)]

$$\phi \leftarrow \phi + \nabla_\phi \left\{ \frac{1}{M} \sum_{k=i-M+1}^i \left( \sum_{l=1}^{L^{(k)}} U_l^\theta(\pi^\phi, k) + \gamma \sum_{l=0}^{L^{(k)}} E_l^1(\pi^\phi, k) \right) \right\}$$

25:     **end if**

26: **end for**

---

Let  $v(t, x) = e^{-\beta t} \psi^*(t, x)$  for all  $(t, x) \in [0, T] \times \mathcal{X}$ , then  $v$  is continuously differentiable in  $t$ , i.e.  $v \in C^{1,0}([0, T] \times \mathcal{X})$ , and satisfies (7). Thus, we establish the existence of the solution to (7).

Suppose  $v^* \in C^{1,0}([0, T] \times \mathcal{X})$  is a solution to (7). One can readily see that the transition rates of the Markov process  $\tilde{X}^\pi$ , introduced in Section EC.2, at time  $s$  is given by

$$\sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} q(y | s, \tilde{X}_{s-}^\pi, S) \pi(S | s, \tilde{X}_{s-}^\pi), \quad y \neq \tilde{X}_{s-}^\pi.$$

It follows from Theorem 3.1 in Guo et al. (2015) that we have the Dynkin's formula:

$$\begin{aligned} & \mathbb{E} \left[ v^*(T, \tilde{X}_T^\pi) \mid \tilde{X}_t^\pi = x \right] - v^*(t, x) \\ &= \mathbb{E} \left[ \int_t^T \left( \frac{\partial v^*}{\partial s}(s, \tilde{X}_{s-}^\pi) + \sum_{y \in \mathcal{X}} v^*(s, y) \left[ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} q(y | s, \tilde{X}_{s-}^\pi, S) \pi(S | s, \tilde{X}_{s-}^\pi) \right] \right) ds \mid \tilde{X}_t^\pi = x \right]. \end{aligned}$$

Since  $v^*$  satisfies Equation (7) with  $v^*(T, x) = 0$ , we then infer that

$$\begin{aligned} v^*(t, x) &= \mathbb{E} \left[ \int_t^T \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S | s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, \tilde{X}_{s-}^\pi)) \right\} ds \mid \tilde{X}_t^\pi = x \right] \\ &= J(t, x; \pi), \quad (t, x) \in [0, T] \times \mathcal{X}. \end{aligned}$$

The proof is therefore complete.

*Proof of the optimal stochastic policy (4) in Section 2.2* Let the space  $B([0, T] \times \mathcal{X})$  and the endowed norm  $\|\cdot\|$  be as specified in the proof of Lemma 1. We then define an operator  $\bar{\mathcal{L}}$  on  $B([0, T] \times \mathcal{X})$  by

$$\bar{\mathcal{L}}\psi(t, x) := e^{\beta t} \int_t^T \sup_{\pi \in \Pi} \left\{ R^\pi(s, x) + \sum_{S \in \mathcal{A}(x)} \left( \sum_{y \in \mathcal{X}} e^{-\beta s} \psi(s, y) q(y | s, x, S) \right) \pi(S | s, x) \right\} ds,$$

where  $\beta := 2\lambda + 1$ . Using a similar argument as in the proof of Lemma 1, we can identify  $\bar{\mathcal{L}}$  as a contraction operator on the Banach space  $B([0, T] \times \mathcal{X})$  and then establish the existence of the solution in the space  $C^{1,0}([0, T] \times \mathcal{X})$  to the following exploratory HJB equation:

$$\begin{cases} \frac{\partial v}{\partial t}(t, x) + \sup_{\pi \in \Pi} \left\{ \sum_{S \in \mathcal{A}(x)} [H(t, x, S, v(\cdot, \cdot)) - \gamma \log \pi(S | t, x)] \pi(S | t, x) \right\} = 0, & (t, x) \in [0, T] \times \mathcal{X} \\ v(T, x) = 0, & x \in \mathcal{X}. \end{cases} \quad (\text{EC.4})$$

Assume  $v^* \in C^{1,0}([0, T] \times \mathcal{X})$  is a solution to (EC.4), then  $v^*(T, x) = 0$  for all  $x \in \mathcal{X}$  and the following inequality holds for all  $\pi \in \Pi$  and  $(t, x) \in [0, T] \times \mathcal{X}$ :

$$\frac{\partial v^*}{\partial t}(t, x) + \sum_{S \in \mathcal{A}(x)} \{H(t, x, S, v^*(\cdot, \cdot)) - \gamma \log \pi(S | t, x)\} \pi(S | t, x) \leq 0. \quad (\text{EC.5})$$

Then, it follows from Theorem 3.1 in Guo et al. (2015) that

$$\begin{aligned}
& v^*(t, x) \\
&= \mathbb{E} \left[ v^*(T, \tilde{X}_T^\pi) \mid \tilde{X}_t^\pi = x \right] \\
&\quad - \mathbb{E} \left[ \int_t^T \left( \frac{\partial v^*}{\partial s}(s, \tilde{X}_{s-}^\pi) + \sum_{y \in \mathcal{X}} v^*(s, y) \left[ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} q(y \mid s, \tilde{X}_{s-}^\pi, S) \pi(S \mid s, \tilde{X}_{s-}^\pi) \right] \right) ds \mid \tilde{X}_t^\pi = x \right] \\
&\geq \mathbb{E} \left[ \int_t^T \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} \{r(S) - \gamma \log \pi(S \mid t, \tilde{X}_{s-}^\pi)\} \pi(S \mid s, \tilde{X}_{s-}^\pi) ds \mid \tilde{X}_t^\pi = x \right] \\
&= J(t, x; \pi), \quad \text{for all } \pi \in \Pi, (t, x) \in [0, T] \times \mathcal{X}.
\end{aligned}$$

Moreover, the equality in (EC.5) holds for policy  $\pi^*$  defined as follows:

$$\pi^*(S \mid t, x) = \frac{\exp\{\frac{1}{\gamma} H(t, x, S, v^*(\cdot, \cdot))\}}{\sum_{S \in \mathcal{A}(x)} \exp\{\frac{1}{\gamma} H(t, x, S, v^*(\cdot, \cdot))\}}, \quad \text{for } (t, x, S) \in \mathbb{K}.$$

Thus, we conclude that  $v^*(t, x) = J(t, x; \pi^*) = J^*(t, x)$  for all  $(t, x) \in [0, T] \times \mathcal{X}$ . This establishes the characterization of the optimal policy in (4).

*Proof of Proposition 1* Denote

$$\tilde{M}_t := J(t, \tilde{X}_t^\pi; \pi) + \int_0^t \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S \mid s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot \mid s, \tilde{X}_{s-}^\pi)) \right\} ds, \quad t \in [0, T].$$

We then show that  $\{\tilde{M}_t : t \in [0, T]\}$  is an  $\{\mathcal{F}_t^{\tilde{X}^\pi}\}_{t \geq 0}$ -martingale. We integrate the expression for  $J(t, x; \pi)$  in (6), yielding

$$\begin{aligned}
\tilde{M}_t &= \mathbb{E} \left[ \int_t^T \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S \mid s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot \mid s, \tilde{X}_{s-}^\pi)) \right\} ds \mid \tilde{X}_t^\pi \right] \\
&\quad + \left[ \int_0^t \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S \mid s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot \mid s, \tilde{X}_{s-}^\pi)) \right\} ds \right].
\end{aligned} \tag{EC.6}$$

Due to the Markov property of  $\{\tilde{X}_t^\pi : t \in [0, T]\}$ , and also note that the second term on the right side of equation (EC.6) is  $\mathcal{F}_t^{\tilde{X}^\pi}$ -measurable, we obtain

$$\tilde{M}_t = \mathbb{E} \left[ \int_0^T \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S \mid s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot \mid s, \tilde{X}_{s-}^\pi)) \right\} ds \mid \mathcal{F}_t^{\tilde{X}^\pi} \right].$$

Since  $J(T, x; \pi) = 0$  for all  $x \in \mathcal{X}$ , we conclude that  $\tilde{M}_t = \mathbb{E}[\tilde{M}_T \mid \mathcal{F}_t^{\tilde{X}^\pi}]$ , which implies  $\{\tilde{M}_t : t \in [0, T]\}$  is an  $\{\mathcal{F}_t^{\tilde{X}^\pi}\}_{t \geq 0}$ -martingale. Moreover, since  $\tilde{M}_t$  is bounded for any fixed time  $t \in [0, T]$ , it follows that  $\{\tilde{M}_t : t \in [0, T]\}$  is square-integrable.

Conversely, assume that  $\{\tilde{M}_t^\nu : t \in [0, T]\}$  is an  $\{\mathcal{F}_t^{\tilde{X}^\pi}\}_{t \geq 0}$ -martingale and  $v(T, x) = 0$  for all  $x \in \mathcal{X}$ . From this, it follows directly that  $\tilde{M}_t^\nu = \mathbb{E}[\tilde{M}_T^\nu | \mathcal{F}_t^{\tilde{X}^\pi}]$  for all  $t \in [0, T]$ , which is equivalent to

$$\begin{aligned} v(t, \tilde{X}_t^\pi) = & \mathbb{E} \left[ \int_0^T \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S | s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, \tilde{X}_{s-}^\pi)) \right\} ds \mid \mathcal{F}_t^{\tilde{X}^\pi} \right] \\ & - \left[ \int_0^t \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S | s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, \tilde{X}_{s-}^\pi)) \right\} ds \right], \quad t \in [0, T]. \end{aligned} \quad (\text{EC.7})$$

Since the second term on the right side of Equation (EC.7) is  $\mathcal{F}_t^{\tilde{X}^\pi}$ -measurable, then it follows from the Markov property of  $\{\tilde{X}_t^\pi : t \in [0, T]\}$  that

$$\begin{aligned} v(t, \tilde{X}_t^\pi) &= \mathbb{E} \left[ \int_t^T \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S | s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, \tilde{X}_{s-}^\pi)) \right\} ds \mid \mathcal{F}_t^{\tilde{X}^\pi} \right] \\ &= \mathbb{E} \left[ \int_t^T \left\{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} r(S) \pi(S | s, \tilde{X}_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, \tilde{X}_{s-}^\pi)) \right\} ds \mid \tilde{X}_t^\pi \right] \\ &= J(t, \tilde{X}_t^\pi; \pi), \quad t \in [0, T]. \end{aligned}$$

Therefore, we conclude that  $v(t, x) = J(t, x; \pi)$  for all  $(t, x) \in [0, T] \times \mathcal{X}$ .

*Proof of Theorem 1* Denote  $U_t := \int_{(t, T]} p^\top dN_s^\pi - \int_t^T \sum_{S \in \mathcal{A}(X_{s-}^\pi)} r(S) \pi(S | s, X_{s-}^\pi) ds$ , then we have

$$\mathbb{E}[U_t | X_t^\pi] = 0. \quad (\text{EC.8})$$

We further denote  $M_t := J(t, X_t^\pi; \pi) + \int_0^t \{ \sum_{S \in \mathcal{A}(X_{s-}^\pi)} r(S) \pi(S | s, X_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, X_{s-}^\pi)) \} ds$  and  $M_t^\theta := J^\theta(t, X_t^\pi) + \int_0^t \{ \sum_{S \in \mathcal{A}(X_{s-}^\pi)} r(S) \pi(S | s, X_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, X_{s-}^\pi)) \} ds$ . It follows from  $J(T, x; \pi) \equiv 0$  and (EC.8) that

$$\begin{aligned} 2L(\theta) &= \mathbb{E} \left[ \int_0^T (U_t + M_T - M_t^\theta)^2 dt \right] \\ &= \mathbb{E} \left[ \int_0^T [U_t^2 + 2U_t(M_T - M_t^\theta) + (M_T - M_t^\theta)^2] dt \right] \\ &= \mathbb{E} \left[ \int_0^T \left[ U_t^2 + 2U_t \left( \int_t^T \left\{ \sum_{S \in \mathcal{A}(X_{s-}^\pi)} r(S) \pi(S | s, X_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, X_{s-}^\pi)) \right\} ds \right) \right] dt \right] \\ &\quad - \mathbb{E} \left[ \int_0^T 2J^\theta(t, X_t^\pi) \mathbb{E}[U_t | X_t^\pi] dt \right] + \mathbb{E} \left[ \int_0^T (M_T - M_t^\theta)^2 dt \right] \\ &= \mathbb{E} \left[ \int_0^T \left[ U_t^2 + 2U_t \left( \int_t^T \left\{ \sum_{S \in \mathcal{A}(X_{s-}^\pi)} r(S) \pi(S | s, X_{s-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | s, X_{s-}^\pi)) \right\} ds \right) \right] dt \right] \\ &\quad + \mathbb{E} \left[ \int_0^T (M_T - M_t^\theta)^2 dt \right]. \end{aligned}$$

Since the first term does not rely on  $\theta$ , we have

$$\arg \min_{\theta} L(\theta) = \arg \min_{\theta} \mathbb{E} \left[ \int_0^T (M_T - M_t^{\theta})^2 dt \right]. \quad (\text{EC.9})$$

Next, we denote  $\tilde{M}_t := J(t, \tilde{X}_t^{\pi}; \pi) + \int_0^t \{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^{\pi})} r(S) \pi(S | s, \tilde{X}_{s-}^{\pi}) + \gamma \mathcal{H}(\pi(\cdot | s, \tilde{X}_{s-}^{\pi})) \} ds$  and  $\tilde{M}_t^{\theta} := J^{\theta}(t, \tilde{X}_t^{\pi}) + \int_0^t \{ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^{\pi})} r(S) \pi(S | s, \tilde{X}_{s-}^{\pi}) + \gamma \mathcal{H}(\pi(\cdot | s, \tilde{X}_{s-}^{\pi})) \} ds$ . From Proposition 1 we know that  $\{\tilde{M}_t : t \in [0, T]\}$  is an  $\{\mathcal{F}_t^{\tilde{X}^{\pi}}\}_{t \geq 0}$ -martingale. Therefore,

$$\begin{aligned} \mathbb{E} \left[ \int_0^T (\tilde{M}_T - \tilde{M}_t^{\theta})^2 dt \right] &= \mathbb{E} \left[ \int_0^T (\tilde{M}_T - \tilde{M}_t + \tilde{M}_t - \tilde{M}_t^{\theta})^2 dt \right] \\ &= \mathbb{E} \left[ \int_0^T [(\tilde{M}_T - \tilde{M}_t)^2 + (\tilde{M}_t - \tilde{M}_t^{\theta})^2 + 2(\tilde{M}_T - \tilde{M}_t)(\tilde{M}_t - \tilde{M}_t^{\theta})] dt \right] \\ &= \mathbb{E} \left[ \int_0^T (\tilde{M}_T - \tilde{M}_t)^2 dt \right] + \mathbb{E} \left[ \int_0^T (\tilde{M}_t - \tilde{M}_t^{\theta})^2 dt \right] \\ &\quad + 2 \int_0^T \mathbb{E} \left[ (\tilde{M}_t - \tilde{M}_t^{\theta}) \cdot \mathbb{E}[\tilde{M}_T - \tilde{M}_t | \mathcal{F}_t^{\tilde{X}^{\pi}}] \right] dt \\ &= \mathbb{E} \left[ \int_0^T (\tilde{M}_T - \tilde{M}_t)^2 dt \right] + \mathbb{E} \left[ \int_0^T |J(t, \tilde{X}_t^{\pi}; \pi) - J^{\theta}(t, \tilde{X}_t^{\pi})|^2 dt \right]. \end{aligned}$$

Noting that the first term does not rely on  $\theta$ , we obtain

$$\arg \min_{\theta} \mathbb{E} \left[ \int_0^T |\tilde{M}_T - \tilde{M}_t^{\theta}|^2 dt \right] = \arg \min_{\theta} \mathbb{E} \left[ \int_0^T |J(t, \tilde{X}_t^{\pi}; \pi) - J^{\theta}(t, \tilde{X}_t^{\pi})|^2 dt \right].$$

Since the distribution of  $\{\tilde{X}_t^{\pi} : t \in [0, T]\}$  is the same as the distribution of  $\{X_t^{\pi} : t \in [0, T]\}$  together with (EC.9), we conclude that  $\arg \min_{\theta} L(\theta) = \arg \min_{\theta} \text{MSVE}(\theta)$ .

*Proof of Theorem 2* Note that a bounded process  $\xi$  with  $\xi_t \in \mathcal{F}_{t-}^{\tilde{X}^{\pi}}$  is also  $\{\mathcal{F}_t\}_{t \geq 0}$ -predictable. On the other hand, for each  $j = 1, \dots, n$ , the process  $\{N_{j,t}^{\pi} - \int_0^t \sum_{S \in \mathcal{A}(X_{t-}^{\pi})} \lambda P_j(S) \pi(S | s, X_{t-}^{\pi}) dt : t \in [0, T]\}$  is an  $\{\mathcal{F}_t\}_{t \geq 0}$ -martingale. Thus we have

$$\mathbb{E} \left[ \int_0^T \xi_t (p^{\top} dN_t^{\pi}) \right] = \mathbb{E} \left[ \int_0^T \xi_t \left[ \sum_{S \in \mathcal{A}(X_{t-}^{\pi})} r(S) \pi(S | t, X_{t-}^{\pi}) \right] dt \right].$$

It suffices to prove that  $v(t, x) = J(t, x; \pi)$  for all  $(t, x) \in [0, T] \times \mathcal{X}$ , if and only if  $v$  satisfies  $v(T, x) = 0$  for all  $x \in \mathcal{X}$ , and

$$\mathbb{E} \left[ \int_0^T \xi_t \left\{ dv(t, X_t^{\pi}) + \left[ \sum_{S \in \mathcal{A}(X_{t-}^{\pi})} r(S) \pi(S | t, X_{t-}^{\pi}) + \gamma \mathcal{H}(\pi(\cdot | t, X_{t-}^{\pi})) \right] dt \right\} \right] \quad (\text{EC.10})$$

for any bounded process  $\xi$  with  $\xi_t \in \mathcal{F}_{t-}^{\tilde{X}^{\pi}}$ .

We first establish that  $v(t, x) = J(t, x; \pi)$  for all  $(t, x) \in [0, T] \times \mathcal{X}$ , if and only if  $v$  satisfies  $v(T, x) = 0$  for all  $x \in \mathcal{X}$ , and

$$\mathbb{E} \left[ \int_0^T \tilde{\xi}_t \left\{ dv(t, \tilde{X}_t^{\pi}) + \left[ \sum_{S \in \mathcal{A}(\tilde{X}_{t-}^{\pi})} r(S) \pi(S | t, \tilde{X}_{t-}^{\pi}) + \gamma \mathcal{H}(\pi(\cdot | t, \tilde{X}_{t-}^{\pi})) \right] dt \right\} \right] = 0, \quad (\text{EC.11})$$

for any bounded process  $\tilde{\xi}$  with  $\tilde{\xi}_t \in \mathcal{F}_{t-}^{\tilde{X}^\pi}$ . The “only if” part follows immediately from Proposition 1. To establish the “if” part, assume that  $v(T, x) = 0$  for all  $x \in \mathcal{X}$ , and (EC.11) holds for any bounded process  $\tilde{\xi}$  with  $\tilde{\xi}_t \in \mathcal{F}_{t-}^{\tilde{X}^\pi}$ . It follows from Theorem 3.1 of Guo et al. (2015) that the process

$$v(t, \tilde{X}_t^\pi) - \int_0^t \left( \frac{\partial v}{\partial s}(s, \tilde{X}_{s-}^\pi) + \sum_{y \in \mathcal{X}} v(s, y) \left[ \sum_{S \in \mathcal{A}(\tilde{X}_{s-}^\pi)} q(y | s, \tilde{X}_{s-}^\pi, S) \pi(S | s, \tilde{X}_{s-}^\pi) \right] \right) ds$$

defines an  $\{\mathcal{F}_t^{\tilde{X}^\pi}\}_{t \geq 0}$ -martingale, and its boundedness at any fixed time  $t \in [0, T]$  ensures square-integrability.

Then, it follows that

$$\mathbb{E} \left[ \int_0^T \tilde{\xi}_t \left( dv(t, \tilde{X}_t^\pi) - \left( \frac{\partial v}{\partial t}(t, \tilde{X}_t^\pi) + \sum_{y \in \mathcal{X}} v(t, y) \left[ \sum_{S \in \mathcal{A}(\tilde{X}_t^\pi)} q(y | t, \tilde{X}_t^\pi, S) \pi(S | t, \tilde{X}_t^\pi) \right] \right) dt \right) \right] = 0, \quad (\text{EC.12})$$

for any bounded process  $\tilde{\xi}$  with  $\tilde{\xi}_t \in \mathcal{F}_{t-}^{\tilde{X}^\pi}$ . By taking the difference between equations (EC.11) and (EC.12), we obtain that

$$\mathbb{E} \left[ \int_0^T \tilde{\xi}_t \left( \frac{\partial v}{\partial t}(t, \tilde{X}_t^\pi) + \sum_{S \in \mathcal{A}(\tilde{X}_t^\pi)} H(t, \tilde{X}_t^\pi, S, v(\cdot, \cdot)) \pi(S | t, \tilde{X}_t^\pi) + \gamma \mathcal{H}(\pi(\cdot | t, \tilde{X}_t^\pi)) \right) dt \right] = 0 \quad (\text{EC.13})$$

holds for any bounded process  $\tilde{\xi}$  with  $\tilde{\xi}_t \in \mathcal{F}_{t-}^{\tilde{X}^\pi}$ . Define the test function  $\tilde{\xi}_t$  as follows:

$$\tilde{\xi}_t = \text{sgn} \left( \frac{\partial v}{\partial t}(t, \tilde{X}_t^\pi) + \sum_{S \in \mathcal{A}(\tilde{X}_t^\pi)} H(t, \tilde{X}_t^\pi, S, v(\cdot, \cdot)) \pi(S | t, \tilde{X}_t^\pi) + \gamma \mathcal{H}(\pi(\cdot | t, \tilde{X}_t^\pi)) \right),$$

where  $\text{sgn}(\cdot)$  is the sign function. Then, (EC.13) implies that

$$\frac{\partial v}{\partial t}(t, \tilde{X}_t^\pi) + \sum_{S \in \mathcal{A}(\tilde{X}_t^\pi)} H(t, \tilde{X}_t^\pi, S, v(\cdot, \cdot)) \pi(S | t, \tilde{X}_t^\pi) + \gamma \mathcal{H}(\pi(\cdot | t, \tilde{X}_t^\pi)) = 0, \quad t \in [0, T]. \quad (\text{EC.14})$$

It follows from Lemma 1 that  $v(t, x) = J(t, x; \pi)$  for all  $(t, x) \in [0, T] \times \mathcal{X}$ . Hence, we complete the proof of the equivalent condition (EC.11).

Let  $D[0, T]$  denote the space of all mappings  $f : [0, T] \mapsto \mathcal{X}$ . For any process  $Y \in D[0, T]$  and fixed  $t \in [0, T]$ , define the stopped process  $Y_{t-\wedge}^\pi \in D[0, T]$  such that  $Y_{t-\wedge}^\pi(s) = Y_s^\pi$  for  $s \in [0, t)$ , and  $Y_{t-\wedge}^\pi(s) = Y_{t-}^\pi$  for  $s \in [t, T]$ . Note that any process  $\xi$  with  $\xi_t \in \mathcal{F}_{t-}^{\tilde{X}^\pi}$  corresponds to a measurable function  $\xi : [0, T] \times D([0, T]) \mapsto \mathbb{R}$  such that  $\xi_t = \xi(t, X_{t-\wedge}^\pi)$ . Then,  $\xi(t, \tilde{X}_{t-\wedge}^\pi)$  is  $\mathcal{F}_{t-}^{\tilde{X}^\pi}$ -measurable for each  $t \in [0, T]$ . Since the distribution of  $\{X_t^\pi : t \in [0, T]\}$  is the same as the distribution of  $\{\tilde{X}_t^\pi : t \in [0, T]\}$ , it follows that

$$\begin{aligned} & \mathbb{E} \left[ \int_0^T \xi(t, X_{t-\wedge}^\pi) \left\{ dv(t, X_t^\pi) + \left[ \sum_{S \in \mathcal{A}(X_{t-}^\pi)} r(S) \pi(S | t, X_{t-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | t, X_{t-}^\pi)) \right] dt \right\} \right] \\ &= \mathbb{E} \left[ \int_0^T \xi(t, \tilde{X}_{t-\wedge}^\pi) \left\{ dv(t, \tilde{X}_t^\pi) + \left[ \sum_{S \in \mathcal{A}(\tilde{X}_{t-}^\pi)} r(S) \pi(S | t, \tilde{X}_{t-}^\pi) + \gamma \mathcal{H}(\pi(\cdot | t, \tilde{X}_{t-}^\pi)) \right] dt \right\} \right]. \end{aligned} \quad (\text{EC.15})$$

Conversely, any process  $\tilde{\xi}$  with  $\tilde{\xi}_t \in \mathcal{F}_{t-}^{\tilde{X}^\pi}$  corresponds to a measurable function  $\tilde{\xi} : [0, T] \times D([0, T]) \mapsto \mathbb{R}$  such that  $\tilde{\xi}_t = \tilde{\xi}(t, \tilde{X}_{t-\wedge}^\pi)$ . Then,  $\tilde{\xi}(t, X_{t-\wedge}^\pi)$  is  $\mathcal{F}_{t-}^{X^\pi}$ -measurable for each  $t \in [0, T]$ , and equation (EC.15) holds for  $\tilde{\xi}(\cdot, \cdot)$ . This establishes the equivalence between condition (EC.11) and condition (EC.10). The proof is therefore complete.

*Proof of Theorem 3* By the representation (23) of  $g(t, x; \phi) = \nabla_\phi J(t, x; \pi^\phi)$ , we have

$$\begin{aligned} \nabla_\phi J(0, c; \pi^\phi) = & \mathbb{E} \left[ \int_0^T \left( \sum_{S \in \mathcal{A}(X_{t-}^{\pi^\phi})} \nabla_\phi \log \pi^\phi(S | t, X_{t-}^{\pi^\phi}) H(t, X_{t-}^{\pi^\phi}, S, J(\cdot, \cdot; \pi^\phi)) \pi^\phi(S | t, X_{t-}^{\pi^\phi}) \right) dt \right. \\ & \left. + \gamma \int_0^T \nabla_\phi \mathcal{H}(\pi^\phi(\cdot | t, \tilde{X}_{t-}^{\pi^\phi})) dt | X_0^{\pi^\phi} = c \right]. \end{aligned} \quad (\text{EC.16})$$

It suffices to show the first term of (EC.16) also equals to the first term of (24). To simplify the notation, we use  $\mathbb{E}_c$  to denote the expectation conditional on the initial state  $X_0^{\pi^\phi} = c$  in the subsequent derivations.

Denote the arrival time of the  $k$ -th customer as  $T_k$  and the product chosen by the customer as  $\zeta_{T_k}$ , where  $\zeta_{T_k} \in \{0, 1, \dots, n\}$ . For  $j = 1, \dots, n$ , denote

$$I_j(t, x, S) := \nabla_\phi \log \pi^\phi(S | t, x) [J(t, x - A^j; \pi^\phi) - J(t, x; \pi^\phi) + p_j].$$

Then, the first term of (24) can be rewritten as

$$\sum_{k=1}^{\infty} \mathbb{E}_c \left[ \sum_{j=1}^n I_j(T_k, X_{T_k-}^{\pi^\phi}, S_{T_k}^{\pi^\phi}) \mathbf{1}_{\{\zeta_{T_k}=j\}} \mathbf{1}_{\{T_k \leq T\}} \right]. \quad (\text{EC.17})$$

For each  $k \geq 1$ , by taking conditional expectation with respect to  $\sigma(T_k, X_{T_k-}^{\pi^\phi})$  inside the expectation in (EC.17), we obtain

$$\begin{aligned} & \sum_{k=1}^{\infty} \mathbb{E}_c \left[ \sum_{S \in \mathcal{A}(X_{T_k-}^{\pi^\phi})} \left( \sum_{j=1}^n I_j(T_k, X_{T_k-}^{\pi^\phi}, S) P_j(S) \pi^\phi(S | T_k, X_{T_k-}^{\pi^\phi}) \right) \mathbf{1}_{\{T_k \leq T\}} \right] \\ &= \mathbb{E}_c \left[ \int_0^T \sum_{S \in \mathcal{A}(X_{t-}^{\pi^\phi})} \left( \sum_{j=1}^n I_j(t, X_{t-}^{\pi^\phi}, S) P_j(S) \pi^\phi(S | t, X_{t-}^{\pi^\phi}) \right) dN_t^\lambda \right] \\ &= \mathbb{E}_c \left[ \int_0^T \sum_{S \in \mathcal{A}(X_{t-}^{\pi^\phi})} \left( \sum_{j=1}^n I_j(t, X_{t-}^{\pi^\phi}, S) \lambda P_j(S) \pi^\phi(S | t, X_{t-}^{\pi^\phi}) \right) dt \right], \end{aligned} \quad (\text{EC.18})$$

where (EC.18) is derived from the fact that  $\{N_t^\lambda - \lambda t : t \in [0, T]\}$  is an  $\{\mathcal{F}_t\}_{t \geq 0}$ -martingale, and the integrand is  $\{\mathcal{F}_t\}_{t \geq 0}$ -predictable and bounded. The boundedness of the integrand is ensured by Assumption 1 and the inherent boundedness of the state process  $\{X_t^{\pi^\phi} : t \in [0, T]\}$ . On the other hand, given the definition of  $H$  in (2), we have

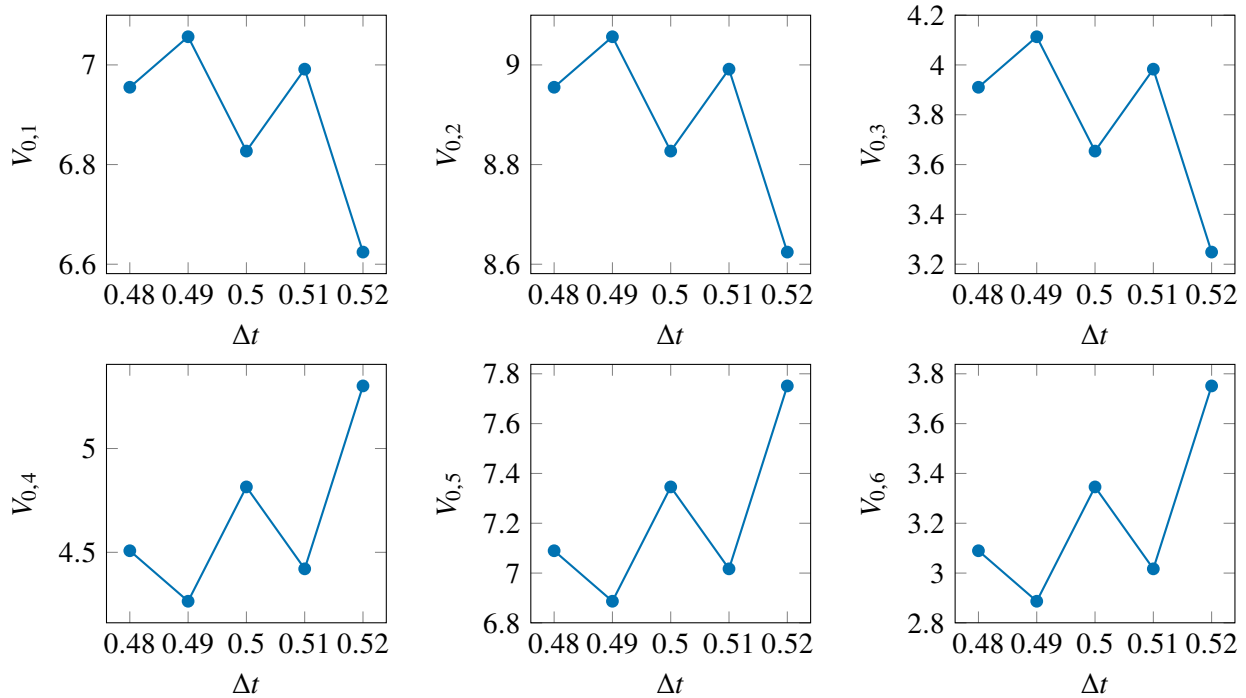
$$H(t, x, S, J(\cdot, \cdot; \pi^\phi)) = \lambda \sum_{j=1}^n p_j P_j(S) + \sum_{y \neq x} \left( \sum_{\{j \in \mathcal{J} : A^j = x - y\}} \lambda P_j(S) \right) J(t, y; \pi^\phi) - \lambda [1 - P_0(S)] J(t, x; \pi^\phi)$$

$$= \sum_{j=1}^n [J(t, x - A^j; \pi^\phi) - J(t, x; \pi^\phi) + p_j] \lambda P_j(S).$$

Thus, (EC.18) is exactly the first term of (EC.16). This concludes the proof.

## EC.5. The instability of ADP time discretization

Figure EC.1 illustrates the variation of the coefficients  $V_{0,i}$  for  $i = 1, \dots, 6$ , selected from the ADP solution in Section 6.3, across five closely spaced levels of time discretization  $\Delta t = 0.48, 0.49, 0.50, 0.51$  and  $0.52$ .



**Figure EC.1** Variations of selected coefficients from the ADP solution across different time discretization levels

## EC.6. Extension to Admission Control in a Single-Server Queue

To demonstrate the generality of our framework, we apply it to another classic problem in Operations Research: admission control in a single-server queue. We assume the customer arrival process and service process are non-homogeneous Poisson process with rate function  $\lambda(t)$  and  $\mu(t)$ , respectively. If  $\mu(t) \equiv \mu$ , that means the service times of customers are i.i.d. exponentially distributed with rate  $\mu$ . An agent is responsible for determining whether to admit arrivals in order to maximize the expected total revenue over a finite horizon  $[0, T]$ . The system has a maximum capacity of  $C$  customers, and it is initially empty. Let  $A_t$  denote the admitted arrival process, and  $D_t$  denote the customer departure process. The queue length (which includes the customer in service), represented by  $X_t = A_t - D_t$ , is the system state. The action space is a binary set  $\{0, 1\}$ , where 0 indicates rejecting an arriving customer, and 1 indicates admitting the customer. Then, a

(randomized) policy  $\pi$  is defined as a mapping from each time-state pair  $(t, x)$  to a Bernoulli distribution with success probability  $p$ . The entropy-regularized value function is given by

$$J(t, x; \pi) := \mathbb{E} \left[ K_1 \int_{(t, T]} dA_s^\pi - K_2 \int_t^T X_{s-}^\pi ds - K_3 X_T^\pi + \gamma \int_t^T \mathcal{H}(\pi(\cdot | s, X_{s-}^\pi)) ds \mid X_t^\pi = x \right].$$

The positive coefficients  $K_1$ ,  $K_2$ ,  $K_3$  are interpreted as follows:  $K_1$  is the jump reward for admitting a customer,  $K_2$  is the holding cost per unit of time for each admitted customer, and  $K_3$  denotes the penalty for an admitted customer not served by the terminal time  $T$ . The earlier analysis can be readily extended to this queueing formulation, with almost identical proof methodologies. We present the results as follows.

For an approximate value function parametrized by  $\theta$ , we define the loss function  $L(\theta)$  as

$$L(\theta) = \frac{1}{2} \mathbb{E} \left[ \int_0^T \left( K_1 \int_{(t, T]} dA_s^\pi - K_2 \int_t^T X_{s-}^\pi ds - K_3 X_T^\pi + \gamma \int_t^T \mathcal{H}(\pi(\cdot | s, X_{s-}^\pi)) ds - J^\theta(t, X_t^\pi) \right)^2 dt \right].$$

It allows us to establish the equivalence theorem, which confirms that the Monte Carlo method for PE remains effective within the queueing framework.

**THEOREM EC.1 (Parallel to Theorem 1).** *It holds that  $\arg \min_\theta L(\theta) = \arg \min_\theta \text{MSVE}(\theta)$ .*

Moreover, we have developed the martingale orthogonality conditions for the queueing model, which form the basis for the TD method in PE.

**THEOREM EC.2 (Parallel to Theorem 2).** *A function  $v \in C^{1,0}([0, T] \times \mathcal{X})$  is the value function associated with the policy  $\pi$ , i.e.  $v(t, x) = J(t, x; \pi)$  for all  $(t, x) \in [0, T] \times \mathcal{X}$ , if and only if it satisfies  $v(T, x) = -K_3 x$  for all  $x \in \mathcal{X}$ , and the following martingale orthogonality condition holds for any bounded process  $\xi$  with  $\xi_t \in \mathcal{F}_{t-}^{X^\pi}$  for all  $t \in [0, T]$ :*

$$\mathbb{E} \left[ \int_0^T \xi_t \left\{ dv(t, X_t^\pi) + K_1 dA_t^\pi - K_2 X_{t-}^\pi dt + \gamma \mathcal{H}(\pi(\cdot | t, X_{t-}^\pi)) dt \right\} \right] = 0.$$

Finally, we derive the formula for PG.

**THEOREM EC.3 (Parallel to Theorem 3).** *Given an admissible parameterized policy  $\pi^\phi$  satisfying Assumption 1, the policy gradient  $\nabla_\phi J(0, c; \pi^\phi)$  admits the following representation:*

$$\begin{aligned} \nabla_\phi J(0, c; \pi^\phi) = \mathbb{E} \left[ \int_{(0, T]} \nabla_\phi \log \pi^\phi(1 | t, X_{t-}^{\pi^\phi}) [J(t, X_{t-}^{\pi^\phi} + 1; \pi^\phi) - J(t, X_{t-}^{\pi^\phi}; \pi^\phi) + K_1] dA_t^{\pi^\phi} \right. \\ \left. + \gamma \int_0^T \nabla_\phi \mathcal{H}(\pi^\phi(\cdot | t, X_{t-}^{\pi^\phi})) dt \right]. \end{aligned}$$

On combining PE with PG, we can now design actor-critic algorithms. We present Algorithm EC.3 as an example. This algorithm employs the neural-network-based approximation for both value and policy functions, and utilizes the Monte Carlo method for PE.

To demonstrate the performance of Algorithm EC.3, we consider a simple example: we set the maximum capacity  $C = 10$ , the time horizon  $T = 20$ , the coefficients  $K_1 = 10$ ,  $K_2 = 1$ ,  $K_3 = 0.1$ . The rate functions for customer arrival process and service process are defined as  $\lambda(t) = 0.5 + 0.3 \cdot \sin(\frac{2\pi t}{T})$  and  $\mu(t) = 0.1 + 0.1 \cdot \frac{t}{T}$ , respectively. However, the RL agent is not aware of these underlying rate functions.

We employ the neural-network-based approximation for value functions and policies, following the 2-NNs approach described in Section 6. We continue to use fully connected networks with each hidden layer followed by a ReLU activation function. The only minor variation in this example is that the outputs of the actor network are one-dimensional, and we directly apply the sigmoid function to convert these outputs into admission probabilities. Algorithm EC.3 is implemented with both the actor and critic networks configured with two hidden layers, each having a width of 8. The hyperparameters are set as follows: batch size  $M = 100$ , entropy factor  $\gamma = 1 \times 10^{-3}$ , learning rate  $\alpha_\theta = 3 \times 10^{-2}$ , and learning rate  $\alpha_\phi = 1 \times 10^{-5}$ . Figure EC.2 illustrates how the average revenues of the updated policies varies throughout the learning process. Upon completing the final update, the simulated average revenues of the resulting policy, achieves a value of 23.930.

We also compare the performance of our RL algorithm with three benchmark policies:

- **Optimal Policy from Discretized Dynamic Programming:** Similar as in Section 6.1.1, for small  $\Delta t$ , solving the DP problem (without entropy regularization) for the discrete-time model can be expected to yield a reliable approximation  $V_{\Delta t}^*(0,0)$  of the true optimal expected revenue  $V^*(0,0)$ . The associated DP equation is given by

$$\begin{cases} V_{\Delta t}^*(t_k, x) = \lambda(t_{k+1})\Delta t \cdot \mathbf{1}_{\{x \leq C-1\}} \cdot \max_{a \in \{0,1\}} \{a \cdot [K_1 + V_{\Delta t}^*(t_{k+1}, x+1) - V_{\Delta t}^*(t_{k+1}, x)]\} + K_2\Delta t \cdot x \\ \quad + \mu(t_{k+1})\Delta t \cdot \mathbf{1}_{\{x \geq 1\}} \cdot [V_{\Delta t}^*(t_{k+1}, x-1) - V_{\Delta t}^*(t_{k+1}, x)] + V_{\Delta t}^*(t_{k+1}, x), \quad \forall k \\ V_{\Delta t}^*(t_K, x) = K_3 \cdot x \end{cases} \quad (\text{EC.19})$$

- **UNIFORM-RANDOM:** At each time and state, the decision to admit or reject an arriving customer is made with an equal probability of 0.5.

- **OPTIMAL-THRESHOLD:** For each state  $\bar{x} \in \{1, \dots, C\}$ , a (deterministic) threshold policy  $\pi^{\bar{x}}$  admits an arriving customer only if the current state is below the threshold  $\bar{x}$ . We execute all these threshold policies and refer to the maximum average revenues obtained as the performance of the OPTIMAL-THRESHOLD policy.

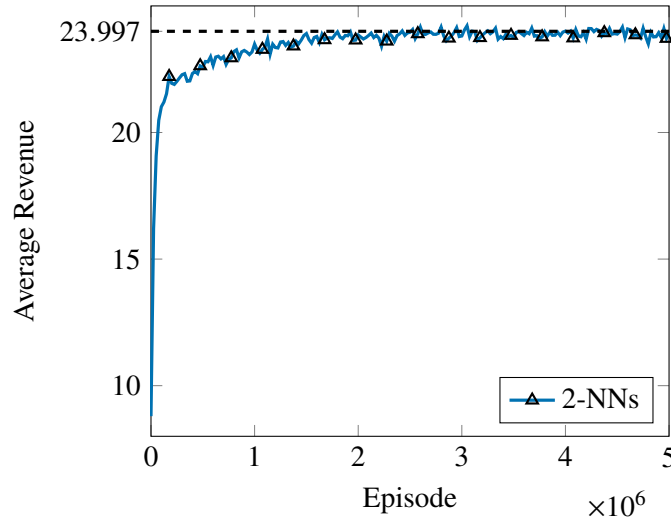
We set  $\Delta t = 0.001$  and solve the DP problem (EC.19), yielding an optimal expected revenue of  $V_{0.001}^*(0,0) = 23.997$ , as indicated in Figure EC.2. Table EC.2 reports the simulated average revenues for our 2-NNs policy and the benchmark policies UNIFORM-RANDOM and OPTIMAL-THRESHOLD. The numerical

**Algorithm EC.3** Actor-Critic Algorithm with Neural Networks for Queueing Control

- 
- 1: **Inputs:** maximum capacity  $C$ , time horizon  $T$ , the values of  $K_1, K_2, K_3$ ; number of episodes  $N$ , batch size  $M$ ; critic network  $J^\theta(t, x)$  and an initial value  $\theta_0$ ; actor network  $\pi^\phi(a | t, x)$  and an initial value  $\phi_0$ ; entropy factor  $\gamma$ , learning rates  $\alpha_\theta, \alpha_\phi$
  - 2: **Required program:** an environment simulator  $(t', x', r') = \text{Environment}(t, x, \pi^\phi(\cdot | \cdot, \cdot); C)$
  - 3: Initialize  $\theta \leftarrow \theta_0, \phi \leftarrow \phi_0$
  - 4: **for**  $i = 1$  to  $N$  **do**
  - 5:   Store  $(\tau_0^{(i)}, x_0^{(i)}) \leftarrow (0, 0)$
  - 6:   Initialize  $l = 0, (t, x) = (0, 0)$
  - 7:   **while** True **do**
  - 8:     Apply  $(t, x)$  and policy  $\pi^\phi$  to the environment simulator to get the next jump time  $t'$ , post-jump state  $x'$  and the jump reward  $r'$   $\triangleright r' = K_1$  for  $x' = x + 1$  and  $r' = 0$  for  $x' = x - 1$
  - 9:     **if**  $t' \geq T$  **then**
  - 10:       Store  $L^{(i)} \leftarrow l, \tau_{L^{(i)}+1}^{(i)} \leftarrow T$
  - 11:       **break**
  - 12:     **end if**
  - 13:     Update  $l \leftarrow l + 1$ .   Store observation at jump time:  $(\tau_l^{(i)}, x_l^{(i)}, r_l^{(i)}) \leftarrow (t', x', r')$
  - 14:     Update  $(t, x) \leftarrow (t', x')$
  - 15:   **end while**
  - 16:   Denote  $D(t_1, t_2, x; \theta) := \int_{t_1}^{t_2} [J^\theta(s, x)]^2 ds, b(t_1, t_2, x; \theta) := \int_{t_1}^{t_2} J^\theta(s, x) ds$
  - 17:    $E(t_1, t_2, x; v, \pi) := \int_{t_1}^{t_2} v(s) \mathcal{H}(\pi(\cdot | s, x)) ds, w(t_1, t_2, x; \theta) := \int_{t_1}^{t_2} s \cdot J^\theta(s, x) ds$
  - 18:   Denote  $D_l(\theta, i) := D(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \theta), b_l(\theta, i) := b(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \theta), w_l(\theta, i) := w(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \theta)$
  - 19:    $\bar{C}_l(\pi^\phi, i) := \sum_{l'=l+1}^{L^{(i)}} [r_{l'}^{(i)} - K_2 x_{l'}^{(i)} (\tau_{l'+1}^{(i)} - \tau_{l'}^{(i)}) + \gamma E(\tau_{l'}^{(i)}, \tau_{l'+1}^{(i)}, x_{l'}^{(i)}; \mathbf{1}, \pi^\phi)]$
  - 20:    $E_l(\theta, \pi^\phi, i) := E(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; b(\tau_l^{(i)}, \cdot, x_l^{(i)}; \theta), \pi^\phi), E_l^1(\pi^\phi, i) = E(\tau_l^{(i)}, \tau_{l+1}^{(i)}, x_l^{(i)}; \mathbf{1}, \pi^\phi)$
  - 21:    $\bar{U}_l^\theta(\pi^\phi, i) := \log \pi^\phi(1 | \tau_l^{(i)}, x_{l-1}^{(i)}) [J^\theta(\tau_l^{(i)}, x_{l-1}^{(i)} + 1) - J^\theta(\tau_l^{(i)}, x_{l-1}^{(i)}) + K_1]$
  - 22:   Store  $\mathcal{I}^{(i)} = \{1 \leq l \leq L^{(i)} : x_l^{(i)} = x_{l-1}^{(i)} + 1\}$   $\triangleright$  Store indices of customer admission transitions
  - 23:   **if**  $i = 0 \pmod{M}$  **then**  $\triangleright$  Perform an update using  $M$  episodes generated under the policy
  - 24:     Update critic network:
  - 25:       Update actor network:
- $$\theta \leftarrow \theta - \alpha_\theta \nabla_\theta \left\{ \frac{1}{M} \sum_{k=i-M+1}^i \sum_{l=0}^{L^{(k)}} \left( \frac{1}{2} D_l(\theta, k) + b_l(\theta, k) [K_2 \cdot x_l^{(k)} \tau_{l+1}^{(k)} + K_3 \cdot x_{L^{(k)}}^{(k)}] - b_l(\theta, k) \bar{C}_l(\pi^\phi, k) - K_2 x_l^{(k)} w_l(\theta, k) - \gamma E_l(\theta, \pi^\phi, k) \right) \right\}$$
- $$\phi \leftarrow \phi + \alpha_\phi \nabla_\phi \left\{ \frac{1}{M} \sum_{k=i-M+1}^i \left( \sum_{l \in \mathcal{I}^{(k)}} \bar{U}_l^\theta(\pi^\phi, k) + \gamma \sum_{l=0}^{L^{(k)}} E_l^1(\pi^\phi, k) \right) \right\}$$
- 26:     **end if**
  - 27: **end for**
-

results demonstrate that our RL algorithm significantly outperforms the UNIFORM-RANDOM and OPTIMAL-THRESHOLD benchmarks, and its performance approaches that of the optimal policy from DP with a negligible gap of only 0.3%.

**Figure EC.2** Average revenues of Algorithm EC.3 over episodes for the example in Section EC.6



**Table EC.2** Simulation results for selected policies in Section EC.6

| Policy            | Average revenue | 99%-CI ( $\pm$ ) | $\frac{\text{Average revenue}}{V_{0.001}^*(0,0)}$ (%) | Time cost (second) |
|-------------------|-----------------|------------------|-------------------------------------------------------|--------------------|
| 2-NNs             | 23.930          | 0.390            | 99.72                                                 | 1,526              |
| UNIFORM-RANDOM    | 8.603           | 0.415            | 35.85                                                 | *                  |
| OPTIMAL-THRESHOLD | 13.358          | 0.325            | 55.67                                                 | *                  |

\* indicates that the method involves neither solving LP problems nor learning; the time cost is virtually zero.