

# Structured Critic Adaptation

Soobin Hwang \*

Complied on August 14, 2026

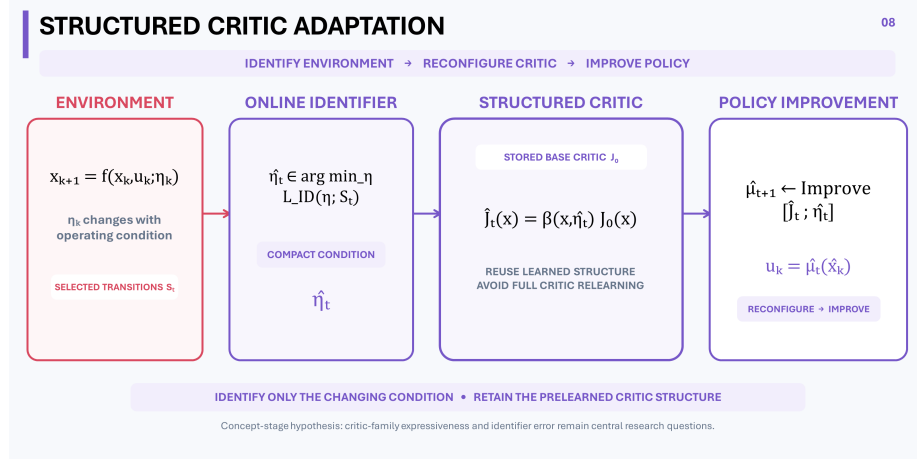


Figure 1: Selected transitions identify the environment, a learned structure re-configures a stored critic, and the resulting critic improves the policy.

Physical interaction is indexed by  $k$ , whereas identification and policy updates are indexed by  $t$ . The environment parameter  $\eta_k$  may vary with physical conditions;  $\hat{\eta}_t$  is the estimate available at learner update  $t$ , and  $k(t)$  is the latest physical sample associated with that update. Identification from a window assumes that  $\eta_k$  is constant or slowly varying over that window, or explicitly models its evolution.

## 1 Problem definition

**Online Learning-Based Optimal Control** requires a critic appropriate to the current decision problem. Consider a family of environments

$$x_{k+1} = f(x_k, u_k; \eta_k) + w_k, \quad y_k = h(x_k) + v_k, \quad (1)$$

where  $x_k$ ,  $u_k$ , and  $y_k$  are the state, control input, and measurement;  $f$  and  $h$  are the state-transition and measurement maps; and  $w_k$  and  $v_k$  represent

\*Soobin Hwang is Integrated M.S./Ph.D. student at KAIST, [tnqls2361@kaist.ac.kr](mailto:tnqls2361@kaist.ac.kr)

process and measurement uncertainty. The compact parameter  $\eta_k$  represents control-relevant changes in dynamics and, when declared, in the stage cost  $g$  or constraints. For notational simplicity, assume the state is available; otherwise, an estimator supplies  $\hat{x}_k$ .

Unless the evolution of  $\eta_k$  is included in a Markov decision state, a condition-specific critic  $J_\eta$  below refers to the problem with  $\eta$  fixed over its value horizon. Material variation over that horizon instead requires a nonstationary formulation such as [Nonstationary Infinite-Horizon OCP](#).

Let the transition record and the set selected at learner update  $t$  be

$$d_k = (x_k, u_k, g_k, x_{k+1}), \quad g_k := g(x_k, u_k; \eta_k), \quad \mathcal{S}_t = \{d_k : k \in \mathcal{I}_t\}, \quad (2)$$

where  $\mathcal{I}_t$  is the selected set of physical transition indices. These recent transitions support online identification:

$$\hat{\eta}_t \in \arg \min_{\eta} \mathcal{L}_{\text{id}}(\eta; \mathcal{S}_t). \quad (3)$$

The loss  $\mathcal{L}_{\text{id}}$  must declare which effects identify  $\eta$ : transition residuals when  $\eta$  changes the dynamics, observed costs when it changes the objective, or constraint/activity information when it changes feasibility. A parameter that affects only an unobserved cost or constraint cannot be identified from state transitions alone. The single  $\hat{\eta}_t$  formulation assumes local constancy or sufficiently slow variation over  $\mathcal{I}_t$ ; otherwise, a trajectory-valued condition model is required.

Let  $J_0$  be a stored approximate base critic for a declared reference condition  $\eta_0$ , and let  $\beta_\omega(x, \eta)$  be a structural map learned over representative conditions. The current critic is reconfigured as

$$\hat{J}_t(x) := \beta_\omega(x, \hat{\eta}_t) J_0(x), \quad (4)$$

and then used for policy improvement:

$$\hat{\mu}_{t+1} \leftarrow \text{Improve}[\hat{J}_t; \hat{\eta}_t]. \quad (5)$$

The generic Improve operator may denote a declared model-based Bellman improvement or a policy-network update. One-step lookahead is therefore a possible implementation, not part of the Theme definition. Here adaptation means reconfiguration through the identified condition, not unrestricted online retraining of all critic parameters. The multiplicative  $\beta_\omega J_0$  form is the current schematic hypothesis; additive corrections or critic-bank interpolation require separate definitions.

## 2 Why this problem matters

Direct adaptive-DP updates can repeatedly modify the full critic whenever the operating condition changes. Such updates may fit the current condition while overwriting behavior useful under earlier conditions, so a recurring condition can require relearning.

If the environment family is organized by an identifiable parameter  $\eta$ , a prelearned relation from  $(x, \eta)$  to critic shape can reuse structure across conditions. Online computation is then concentrated on estimating  $\eta$  and improving the policy from the reconfigured critic. This is a conditional opportunity, not a guaranteed speedup: it depends on identifiability, offline coverage, and the expressiveness of the critic family.

### 3 Key challenges

- **Environment representation and identification:**  $\eta$  must capture control-relevant variation and be identifiable from the limited, policy-dependent data in  $\mathcal{S}_t$ .
- **Structured-family expressiveness:** the family  $\beta_\omega(x, \eta)J_0(x)$  may be too restrictive, especially near zeros or sign changes of  $J_0$ .
- **Error propagation and changing conditions:** identification delay, structural approximation error, stale parameters, and model error perturb both the critic and the improved policy.
- **Bellman and closed-loop consistency:** a plausibly shaped critic is not automatically Bellman-consistent, stabilizing, constraint-satisfying, or reliable outside the representative offline conditions.

### 4 A conventional approach — direct critic adaptation

A conventional adaptive-DP realization updates critic parameters directly from recent data and then improves the policy:

$$\begin{aligned}\theta_{t+1} &\approx \arg \min_{\theta} \mathcal{L}_{\text{critic}}(\theta; \mathcal{S}_t), \\ \hat{\mu}_{t+1} &\leftarrow \text{Improve} \left[ \hat{J}_{\theta_{t+1}} \right].\end{aligned}\tag{6}$$

One representative choice for  $\mathcal{L}_{\text{critic}}$  is a sampled one-step TD loss. Holding the target critic  $\hat{J}_{\hat{\theta}_t}$  fixed during learner update  $t$  and taking  $0 < \alpha < 1$ , define

$$\begin{aligned}y_{k,t}^{\text{TD}} &= g_k + \alpha \hat{J}_{\hat{\theta}_t}(x_{k+1}), \\ \mathcal{L}_{\text{critic}}(\theta; \mathcal{S}_t) &= \frac{1}{|\mathcal{S}_t|} \sum_{k \in \mathcal{I}_t} \left( \hat{J}_{\theta}(x_k) - y_{k,t}^{\text{TD}} \right)^2.\end{aligned}\tag{7}$$

This form evaluates the policy that generated the transitions, subject to the usual on-policy or correction assumptions. If  $\mathcal{S}_t$  mixes behavior policies or materially different environment conditions, the window must be restricted or an explicit off-policy/condition correction must be supplied. If a model and action minimization are available, an empirical Bellman-optimality target can be used instead. The selected critic loss and its policy-evaluation or policy-improvement role must be declared.

This can track the currently visited condition, but it does not explicitly preserve a reusable relation between the environment parameter and critic shape. Under limited data, it may repeatedly overwrite and relearn critic behavior as conditions change. This is a baseline mechanism, not a claim that every direct critic update necessarily forgets.

## 5 Candidate direction — identifier-conditioned critic reconfiguration

Offline, suppose approximate critics  $\{\hat{J}^{(m)}\}$  are available for representative conditions  $\{\eta^{(m)}\}$  under a common state representation, objective, discount, and solution concept. In the optimal-control case,  $\hat{J}^{(m)} \approx J_{\eta^{(m)}}^*$ , where  $J_{\eta}^*$  denotes the optimal critic for the fixed-condition problem. Hold the reference critic  $J_0$  fixed and fit the structural map through

$$\omega^\dagger \in \arg \min_{\omega} \sum_m \left\| \hat{J}^{(m)}(\cdot) - \beta_{\omega}(\cdot, \eta^{(m)}) J_0(\cdot) \right\|_{\nu_m}^2, \quad (8)$$

where  $\nu_m$  is a declared weighting distribution or measure over control-relevant states and  $\|e\|_{\nu_m}^2 := \int |e(x)|^2 d\nu_m(x)$ . If  $J_0$  is to be learned jointly instead, it must appear as an optimization variable and a normalization is needed to remove the scaling ambiguity between  $J_0$  and  $\beta_{\omega}$ . Online operation then follows

$$\begin{aligned} \mathcal{S}_t &\xrightarrow{\text{online identification}} \hat{\eta}_t, \\ \hat{J}_t(x) &= \beta_{\omega^\dagger}(x, \hat{\eta}_t) J_0(x), \\ \hat{\mu}_{t+1} &\leftarrow \text{Improve}[\hat{J}_t; \hat{\eta}_t]. \end{aligned} \quad (9)$$

A useful diagnostic, rather than a guarantee, separates two critic errors. For compactness, let  $\eta_t := \eta_{k(t)}$  denote the physical environment condition associated with learner update  $t$ . The norm below must use the same declared control-relevant domain or measure for all terms:

$$\begin{aligned} \left\| \hat{J}_t - J_{\eta_t}^* \right\| &\leq \left\| \beta_{\omega^\dagger}(\cdot, \hat{\eta}_t) J_0 - \beta_{\omega^\dagger}(\cdot, \eta_t) J_0 \right\| \\ &\quad + \left\| \beta_{\omega^\dagger}(\cdot, \eta_t) J_0 - J_{\eta_t}^* \right\|. \end{aligned} \quad (10)$$

The first term audits sensitivity to identification error; the second audits the structured critic family. Policy-improvement and model errors remain additional. Research questions include how to select  $\eta$ , how to learn  $J_0$  and  $\beta_{\omega}$ , when an additive or critic-bank structure is more appropriate, how to represent uncertainty in  $\hat{\eta}_t$ , and when to reject reconfiguration in favor of a safe fallback.

Future evidence should compare a fixed  $J_0$ , direct critic adaptation, the structured method with oracle  $\eta$ , and the method with estimated  $\eta$ . Relevant metrics include identification error and delay, critic and Bellman error, closed-loop performance, violations, update latency, memory, and return-to-prior-condition performance under abrupt, gradual, recurring, and held-out conditions.

## 6 Application domains

- **Mobility controllers under changing mass, load, road–tire condition, temperature, or component aging:** estimate a compact condition parameter and reconfigure a stored critic so recurring operating regimes can reuse previously learned long-run control structure.

- **Electric-drive and mechatronic control** with compact identifiable plant parameters, such as resistance, inertia, friction, or load, that can condition a reusable critic family and its subsequent policy improvement.
- **Systems that revisit a finite or smoothly parameterized environment family:** recover or interpolate an appropriate critic when a known condition returns, rather than relearning the full value function from the latest trajectory. The benefit depends on identification quality and offline coverage of that environment family.