

CSCI 145 Problem 2: The Optimal Control Variate

July 20, 2026

Submission Instructions

Please upload your work to Gradescope by **11:59pm on DATE**.

- Turn in a single PDF of *handwritten* solutions. Typed (including L^AT_EX) solutions will not be accepted.
- Problems are grouped into three-week **cycles**. Once per cycle, you will meet with the grader and present the solution to *one* problem from that cycle, selected at random when you arrive — so bring your work and be ready to explain any of them.

Problem 2: The Optimal Control Variate

In lecture we halved a Monte Carlo estimator's error without rolling a single extra die: we subtracted off a surrogate whose average we already knew, scaled by a coefficient guessed out of the air ($c = \frac{1}{2}$). Was the guess lucky, and how much can the trick ever buy us? In Problem 1 correlation was the villain that put a floor under a poll; here it answers both questions.

We want $\mu = \mathbb{E}[f(X)]$, and alongside the target $f(X)$ we have a *surrogate* $g(X)$ whose expectation $\mathbb{E}[g(X)]$ we know exactly. In class we built the adjusted Monte Carlo estimator

$$\hat{\mu}_c = \frac{1}{n} \sum_{i=1}^n \left[f(X^{(i)}) - c(g(X^{(i)}) - \mathbb{E}[g(X)]) \right],$$

where $X^{(1)}, \dots, X^{(n)}$ are independent samples distributed like X , and we proved it unbiased for *every* constant c . So all choices of c are correct on average; they differ only in variance.

Part A: The Variance Is a Parabola

Show that

$$\text{Var}(\hat{\mu}_c) = \frac{1}{n} \left(\text{Var}(f(X)) - 2c \text{Cov}(f(X), g(X)) + c^2 \text{Var}(g(X)) \right).$$

(Hint: the samples are independent, so the variance of the average is the variance of one adjusted sample divided by n , the $\rho = 0$ case of Problem 1. Then expand the variance of $f(X) - c g(X)$ with the variance-of-a-sum identity from class; constants shift nothing.) As a function of c , is this parabola smiling or frowning, and why does that matter for what comes next?

Part B: The Bottom of the Parabola

Find the coefficient c^* minimizing $\text{Var}(\hat{\mu}_c)$ by setting the derivative in c to zero. Your answer should be

$$c^* = \frac{\text{Cov}(f(X), g(X))}{\text{Var}(g(X))}.$$

Sanity-check it: what does the formula say when f and g are independent, and does that match what the adjustment *should* do?

Part C: How Much It Saves

Substitute c^* back into Part A and show that

$$\text{Var}(\hat{\mu}_{c^*}) = \frac{\text{Var}(f(X))}{n} (1 - \rho^2), \quad \text{where} \quad \rho = \frac{\text{Cov}(f(X), g(X))}{\sqrt{\text{Var}(f(X)) \text{Var}(g(X))}}$$

is the correlation between target and surrogate. Now apply it to the class example, with f the larger of two dice and g their sum, using the exact moments from enumerating the 36 outcomes:

$$\text{Var}(f(X)) = \frac{2555}{1296}, \quad \text{Cov}(f(X), g(X)) = \frac{35}{12}, \quad \text{Var}(g(X)) = \frac{35}{6}.$$

Compute c^* and the reduction factor $1 - \rho^2$, compare c^* against the value lecture guessed, and say how many games the adjustment saves. Conclude in one sentence: what two properties make a surrogate g worth having, and which of them decides *how much* it is worth?

Where this goes next. Stochastic gradient descent estimates the gradient by Monte Carlo, so everything you proved here applies to it verbatim, and in the final week the variance-minimizing *baseline* an agent subtracts from its returns turns out to be exactly your c^* .