
Log-Derivative Trick

Beier Zhu

In machine learning, we often encounter optimization problems involving expectations over a parameterized distribution:

$$\mathcal{L}(\theta) = \mathbb{E}_{x \sim p_\theta}[f(x)]. \quad (1)$$

Optimizing such objectives is challenging because the sampling operation introduces a non-differentiable stochastic node, preventing gradients from being directly propagated through the sampling process.

In our previous note, *Gumbel-Max*, *Gumbel-Softmax* and *Straight-Through*, we introduced a reparameterization-based approach that enables differentiable sampling for categorical distributions. In this note, we introduce a more general approach, known as the **log-derivative trick** (or score function estimator), which avoids differentiating through the sampling process. Starting from the objective, we have:

$$\nabla_\theta \mathcal{L}(\theta) = \nabla_\theta \mathbb{E}_{x \sim p_\theta}[f(x)] \quad (2)$$

$$= \nabla_\theta \int p_\theta(x) f(x) dx \quad (3)$$

$$= \int f(x) \nabla_\theta p_\theta(x) dx. \quad (4)$$

Using the following identity:

$$\nabla_\theta \log p_\theta(x) = \frac{\nabla_\theta p_\theta(x)}{p_\theta(x)}, \quad (5)$$

Substituting this into the gradient expression gives:

$$\nabla_\theta \mathcal{L}(\theta) = \int p_\theta(x) f(x) \nabla_\theta \log p_\theta(x) dx \quad (6)$$

$$= \mathbb{E}_{x \sim p_\theta}[f(x) \nabla_\theta \log p_\theta(x)]. \quad (7)$$

This transformation is the key idea of the log-derivative trick: instead of differentiating through the sampled variable x , we move the gradient operation onto the log probability of the sampling distribution.

In practice, the expectation can be estimated using Monte Carlo sampling. Given a batch of samples $\mathcal{B} = \{x_1, x_2, \dots, x_n\} \sim p_\theta^n(x)$, the gradient estimator becomes:

$$\nabla_\theta \tilde{\mathcal{L}}(\theta) = \frac{1}{n} \sum_{i=1}^n f(x_i) \nabla_\theta \log p_\theta(x_i). \quad (8)$$

Therefore, a corresponding surrogate objective can be defined as:

$$\tilde{\mathcal{L}}(\theta) = \frac{1}{n} \sum_{i=1}^n f(x_i) \log p_\theta(x_i), \quad (9)$$

whose gradient gives the Monte Carlo estimator above.