

online learning

Authors

Alexander Jung*, Department of Computer Science, Aalto University, Espoo, Finland

*Corresponding author: alex.jung@aalto.fi

Abstract

Online learning refers to machine learning (ML) methods that process data sequentially: at each discrete time instant a new data point arrives and the method updates its current model parameters, in contrast with solving empirical risk minimization (ERM) once over a fixed training set. The canonical update is a gradient step on the loss of the newest data point, which is online gradient descent (online GD). For data points that are realizations of independent and identically distributed (i.i.d.) random variables (RVs), the expected update is a gradient step on the risk, and online GD coincides with stochastic gradient descent (SGD) with batch size one. Performance is measured by the regret, the accumulated loss minus that of the best fixed model parameters in hindsight. With a suitably decaying learning rate, the average regret per round of online GD vanishes as the number of rounds grows.

Definition

A Finnish Meteorological Institute (FMI) weather station delivers one new data point every evening: the day's morning minimum and its maximum daytime temperature. A method that first collects a complete training set never starts forecasting. Online learning instead updates the hypothesis every day, using the newest data point (Jung, 2022, Sect. 4.7).

Online learning refers to machine learning (ML) methods that process data sequentially: data points $\mathbf{z}^{(t)}$ arrive at discrete time instants $t = 1, 2, \dots$, and at each instant the method updates its current model parameters $\mathbf{w}^{(t)}$, or hypothesis $h^{(t)}$, using $\mathbf{z}^{(t)}$ (Fig. 1). This contrasts with solving empirical risk minimization (ERM) once over a fixed training set. The time index is also the iteration index: one new data point means one update.

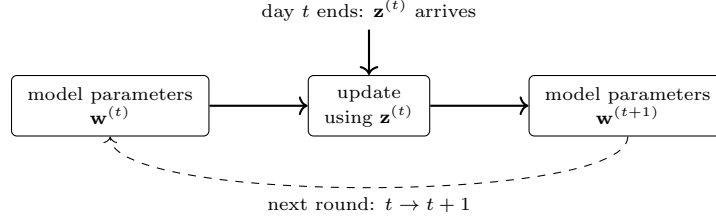


Figure 1: One round of online learning. The model parameters $\mathbf{w}^{(t)}$ are updated as soon as the data point $\mathbf{z}^{(t)}$ of time instant t arrives, and the updated model parameters $\mathbf{w}^{(t+1)}$ enter the next round. In online gradient descent (online GD) the update is a gradient step on $L(\mathbf{z}^{(t)}, \cdot)$

The canonical update is a gradient step on the loss of the newest data point,

$$\mathbf{w}^{(t+1)} = \mathbf{w}^{(t)} - \eta^{(t)} \nabla_{\mathbf{w}} L(\mathbf{z}^{(t)}, \mathbf{w}^{(t)}),$$

which is online GD. The iteration applies no single fixed operator: the update at time t is the gradient step operator of the newest loss component, which changes from round to round, so a fixed point interpretation with one operator is not meaningful here. For data points that are realizations of independent and identically distributed (i.i.d.) random variables (RVs), the expected update is a gradient step on the risk, whose fixed points are the stationary points of the risk; online GD then coincides with stochastic gradient descent (SGD) with batch size one.

The quality of an online method is measured in hindsight. After T rounds, the accumulated loss $\sum_{t=1}^T L(\mathbf{z}^{(t)}, \mathbf{w}^{(t)})$ is compared with the accumulated loss of the best fixed model parameters chosen after all data points have been seen; the difference is the regret. With a suitably decaying learning rate, online GD keeps the average regret per round vanishing as T grows (Hazan, 2022).

For linear regression on the temperature stream, the model parameters that minimize the average squared error loss over all days up to time t need not be recomputed from scratch each day: the solution at time $t + 1$ reuses the matrix computations of time t (Jung, 2022, Sect. 4.7).

Cross References

online GD, online algorithm, SGD, GD, gradient step, ERM, regret, risk, linear regression

References

1. Hazan (2022). Introduction to Online Convex Optimization. 2nd ed., MIT Press, Cambridge, MA, USA.

2. Jung (2022). Machine Learning: The Basics. Springer Nature, Singapore, Singapore.