

Introduction to Machine Learning

Practical Aspects of Training Deep Neural Networks

Prof. Subrahmanya Swamy Peruru
Department of Electrical Engineering
IIT Kanpur

Gradient Descent

$$\mathbf{w}_{opt} = \arg \min_{\mathbf{w}} L(\mathbf{w})$$

- Initialize $\mathbf{w}$ as $\mathbf{w}^{(0)}$
- For iteration $t = 0, 1, 2, \dots$ (or until convergence)
 - Calculate the gradient $\mathbf{g}^{(t)} = \nabla_{\mathbf{w}} L(\mathbf{w}^{(t)})$
 - Set the learning rate η_t
 - Update parameters in the opposite direction of gradient

$$\mathbf{w}^{(t+1)} = \mathbf{w}^{(t)} - \eta_t \mathbf{g}^{(t)}$$

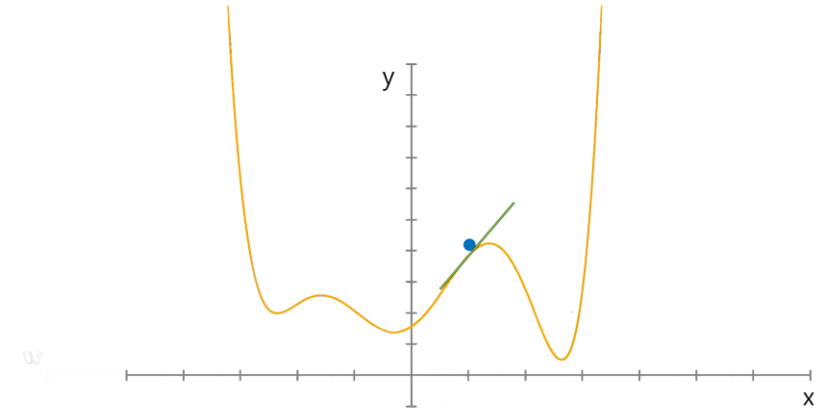


Fig Src: <https://telefonicatech.uk/blog/introduction-to-artificial-neural-networks-part-two-gradient-descent-backpropagation-supervised-unsupervised-learning/>

Stochastic Gradient Descent (SGD)

- Consider a loss function of the form $L(\mathbf{w}) = \frac{1}{N} \sum_{n=1}^N \ell_n(\mathbf{w})$

- The gradient in this case can be written as

$$\mathbf{g} = \nabla_{\mathbf{w}} L(\mathbf{w}) = \nabla_{\mathbf{w}} \left[\frac{1}{N} \sum_{n=1}^N \ell_n(\mathbf{w}) \right] = \frac{1}{N} \sum_{n=1}^N \nabla_{\mathbf{w}} \ell_n(\mathbf{w}) = \frac{1}{N} \sum_{n=1}^N \mathbf{g}_n$$

Expensive to compute – requires doing it for all the training examples in each iteration ☹️

- Stochastic Gradient Descent (SGD) approximates $\mathbf{g}$ using a single training example
- At iter. t , pick an index $i \in \{1, 2, \dots, N\}$ uniformly randomly and approximate $\mathbf{g}$ as

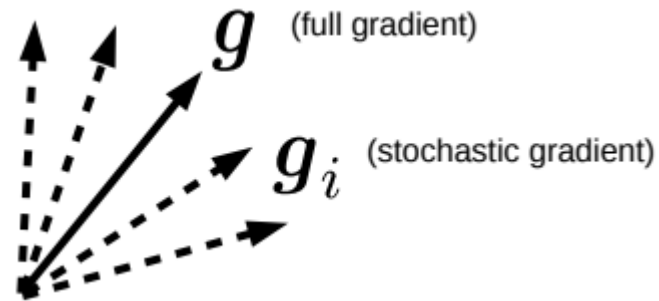
$$\mathbf{g} \approx \mathbf{g}_i = \nabla_{\mathbf{w}} \ell_i(\mathbf{w})$$

Can show that $\mathbf{g}_i$ is an unbiased estimate of $\mathbf{g}$, i.e., $\mathbb{E}[\mathbf{g}_i] = \mathbf{g}$

- May take more iterations than GD to converge but each iteration is much faster 😊
 - SGD per iter cost is $O(D)$ whereas GD per iter cost is $O(ND)$

Minibatch SGD

- Gradient approximation using a single training example may be noisy



The approximation may have a **high variance** – may slow down convergence, updates may be unstable, and may even give sub-optimal solutions (e.g., local minima where GD might have given global minima)

- We can use $B > 1$ unif. rand. chosen train. ex. with indices $\{i_1, i_2, \dots, i_B\} \in \{1, 2, \dots, N\}$
- Using this “minibatch” of examples, we can compute a minibatch gradient

$$\mathbf{g} \approx \frac{1}{B} \sum_{b=1}^B \mathbf{g}_{i_b}$$

- Averaging helps in reducing the variance in the stochastic gradient
- Time complexity is $O(BD)$ per iteration in this case

Training with Mini-batch SGD

Training Procedure

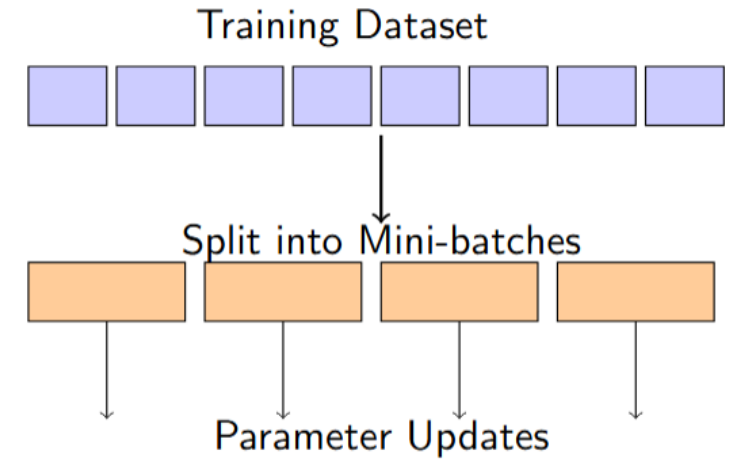
1. Initialize parameters $\mathbf{w}$
2. For **epoch**:
 1. Randomly shuffle the training dataset
 2. Split the dataset into mini-batches of size B
 3. For each mini-batch $b = 1, 2 \dots \frac{N}{B}$

- Compute gradient on the b^{th} mini-batch $\mathcal{B}_b$

$$\mathbf{g} = \frac{1}{B} \sum_{i \in \mathcal{B}_b} \nabla_{\mathbf{w}} \ell_i(\mathbf{w})$$

- Update parameters

$$\mathbf{w} \leftarrow \mathbf{w} - \eta \mathbf{g}$$



One pass through all batches = **1 Epoch**

If dataset size = N and batch size = B ,
Updates per epoch = $\frac{N}{B}$

Example: $N = 8, B = 2$

Note: Shuffling ensures batches are different in every epoch

Adaptive Gradient Methods: AdaGrad

Instead of using a common learning rate for all dimensions $d = 1, 2, \dots, D$

$$w_d^{(t+1)} = w_d^{(t)} - \eta_t g_d^{(t)}$$

Use a separate learning rate for each parameter

$$w_d^{(t+1)} = w_d^{(t)} - e_d^{(t)} g_d^{(t)}$$

AdaGrad:

If some dimension had big updates (indicated by large gradient values in the past), slow down along those directions by using smaller learning rates.

$$e_d^{(t)} = \frac{1}{\sqrt{\epsilon + \sum_{\tau=1}^t \left(g_d^{(\tau)}\right)^2}}$$

Adaptive Gradient Methods: Momentum

7

Can use a momentum term to stabilize gradients by reusing info from past grads

- Move faster along directions where gradients have been consistent
- Slow down along directions where gradient has changed abruptly

The “momentum”
term. Set to 0 at
initialization

β usually set
as 0.9

$$\begin{aligned}\mathbf{m}^{(t)} &= \beta \mathbf{m}^{(t-1)} + \mathbf{g}^{(t)} \\ \mathbf{w}^{(t+1)} &\leftarrow \mathbf{w}^{(t)} - \eta_t \mathbf{m}^{(t)}\end{aligned}$$

Some advanced methods combine momentum and AdaGrad type techniques

- RMS-Prop
- Adam

Initialization of parameters

Why Initialization Matters: Training with gradient descent is sensitive to the starting weights w .

Poor initialization can lead to:

- **Vanishing gradients:** gradients become extremely small
- **Exploding gradients:** gradients grow very large
- **Symmetry problem:** neurons learn identical features

Common Initialization Schemes

Xavier Initialization (Glorot) – for sigmoid / tanh

$$w_{ij} \sim \mathcal{N}\left(0, \frac{1}{n_{\text{in}}}\right)$$

He Initialization – for ReLU networks

$$w_{ij} \sim \mathcal{N}\left(0, \frac{2}{n_{\text{in}}}\right)$$

n_{in} : number of input connections (fan-in)

Batch Normalization

- Standardize = activation $h_{nk}^{(\ell)}$ should have zero mean and unit variance across all n
- It is achieved by inserting a “batch normalization” layer after each hidden layer
- To do so, during training, (omitting layer number ℓ) we replace each $\mathbf{h}_n$ by $\tilde{\mathbf{h}}_n$

$$\mu_{\mathcal{B}} = \frac{1}{|\mathcal{B}|} \sum_{\mathbf{h} \in \mathcal{B}} \mathbf{h} \quad \sigma_{\mathcal{B}}^2 = \frac{1}{|\mathcal{B}|} \sum_{\mathbf{h} \in \mathcal{B}} (\mathbf{h} - \mu_{\mathcal{B}})^2 \quad \hat{\mathbf{h}}_n = \frac{\mathbf{h}_n - \mu_{\mathcal{B}}}{\sqrt{\sigma_{\mathcal{B}}^2 + \epsilon}}$$

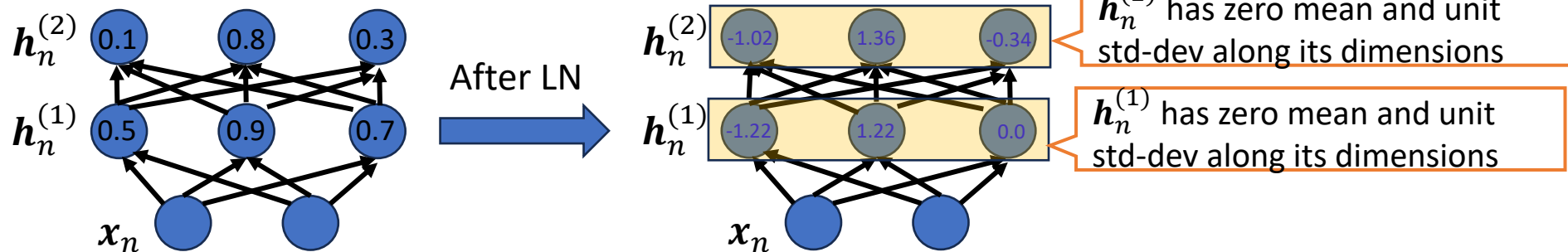
γ and β are trainable
batch-norm parameters

$$\tilde{\mathbf{h}}_n = \gamma \odot \hat{\mathbf{h}}_n + \beta$$

BatchNorm normalizes the signal to stabilize training, but the network is then allowed to learn the best scale and shift again to ensure flexibility

Layer Normalization

- Unlike batch normalization (BN), which we already saw, layer normalization (LN) normalizes each $\mathbf{h}_n$ **across its dimensions** (not across all minibatch examples)
 - Often used for **sequence data models** where BN is difficult to apply
 - Also useful when batch sizes are small where BN statistics (mean/var) aren't reliable
- For an MLP, the LN operation would look like this



Thank you