

MIT 9.520/6.7910
Statistical Learning Theory and Applications

Class 05: Implicit Regularization

Lorenzo Rosasco

Learning algorithm design so far

Pick a representation, then:

- ▶ ERM, penalized/constrained,

$$\min_{w \in \mathbb{R}^d} \underbrace{\frac{1}{n} \sum_{i=1}^n \ell(y_i, w^\top x_i) + \lambda \|w\|^2}_{\hat{L}^\lambda(w)}.$$

- ▶ Optimization by GD/SGD,

$$w_{t+1} = w_t - \gamma \nabla \hat{L}^\lambda(w_t).$$

Beyond ERM

- ▶ Are there other algorithm design principles?
- ▶ So far statistics/regularization separate from computations.

Today we will see how *optimization regularizes implicitly*.

Least squares recap: minimal norm

Recall that for least squares

$$\hat{X}w = \hat{Y}$$

$$\underbrace{\min_{w \in \mathbb{R}^d} \frac{1}{n} \left\| \hat{Y} - \hat{X}w \right\|^2}_{n > d}$$

$$\underbrace{\min_{w \in \mathbb{R}^d} \|w\|, \quad \text{subj. to } \hat{X}w = \hat{Y}}_{n < d}$$

$$\Rightarrow \hat{w}^\dagger = \hat{X}^\dagger \hat{Y}.$$

Least squares recap: pernalized ERM

$$\min_{w \in \mathbb{R}^d} \frac{1}{n} \left\| \hat{Y} - \hat{X} w \right\|^2 + \lambda \|w\|^2$$

$$\Rightarrow \hat{w}_\lambda = (\hat{X}^\top \hat{X} + \lambda n I)^{-1} \hat{X}^\top \hat{Y} = \hat{X}^\top (\hat{X} \hat{X}^\top + \lambda n I)^{-1} \hat{Y}$$

$$\lim_{\lambda \rightarrow 0} (\hat{X}^\top \hat{X} + \lambda n I)^{-1} \hat{X}^\top = \hat{X}^\dagger \quad \Rightarrow \quad \lim_{\lambda \rightarrow 0} \hat{w}_\lambda = \hat{w}^\dagger$$

Outline

The implicit bias of gradient descent

Gradient descent stability and early stopping

Least squares learning with gradient descent

We want to understand the learning properties of

$$\hat{w}_{t+1} = \hat{w}_t - \gamma \frac{1}{n} \hat{X}^\top (\hat{X} \hat{w}_t - \hat{Y}),$$

the gradient descent iteration for the empirical risk,

$$\hat{L}(w) = \frac{1}{2n} \left\| \hat{Y} - \hat{X} w \right\|^2.$$

No penalties no constraints (!): what can we hope for?

Implicit bias of gradient descent

We know that, for suitable γ

$$\lim_{t \rightarrow \infty} \hat{w}_t = \arg \min \hat{L}(w).$$

In fact, we can show that

$$\lim_{t \rightarrow \infty} \hat{w}_t = \hat{w}^\dagger.$$

Understanding the implicit bias of GD for LS

$$\hat{w}_{t+1} = \hat{w}_t - \gamma \frac{2}{n} \hat{X}^\top (\hat{X} \hat{w}_t - \hat{Y}),$$

If $w_0 \in \text{Range}(\hat{X}^\top)$ (e.g. $w_0 = 0$), then $\hat{w}_t \in \text{Range}(\hat{X}^\top)$ for all t .
 $\Rightarrow$ GD converges to a solution in $\text{Range}(\hat{X}^\top)$.

¹Recall that $\text{Range}(\hat{X}^\top) = \text{Null}(\hat{X})^\perp$

Understanding the implicit bias of GD for LS

$$\hat{w}_{t+1} = \hat{w}_t - \gamma \frac{2}{n} \hat{X}^\top (\hat{X} \hat{w}_t - \hat{Y}),$$

If $w_0 \in \text{Range}(\hat{X}^\top)$ (e.g. $w_0 = 0$), then $\hat{w}_t \in \text{Range}(\hat{X}^\top)$ for all t .
 $\Rightarrow$ GD converges to a solution in $\text{Range}(\hat{X}^\top)$.

The minimal norm solution $\hat{w}^\dagger$ is the unique solution in¹ $\text{Range}(\hat{X}^\top)$.

Then,

$$\lim_{t \rightarrow \infty} \hat{w}_t = \hat{w}^\dagger.$$

¹Recall that $\text{Range}(\hat{X}^\top) = \text{Null}(\hat{X})^\perp$

Implicit bias/regularization

Compare

$$\lim_{t \rightarrow \infty} \hat{w}_t = \hat{w}^\dagger,$$

to

$$\lim_{\lambda \rightarrow 0} \hat{w}^\lambda = \hat{w}^\dagger.$$

- ▶ Gradient descent explores solutions with a *bias* towards small norms.
- ▶ The bias is *implicit* in the sense that there are no explicit constraint/penalties.

Implicit bias for other loss functions

Regression

Consider $\ell(x^\top w - y)$ and

$$\hat{w}_{t+1} = \hat{w}_t - \gamma_t \frac{1}{n} \sum_{i=1}^n x_i \ell'(x_i^\top w - y_i),$$

Then $w_t \in \text{Range}(\hat{X}^\top)$, and if $\ell(0) = 0$ and $n > d$ then $\hat{w}_t \rightarrow \hat{w}^\dagger$ (why?).

Implicit bias for other loss functions

Regression

Consider $\ell(x^\top w - y)$ and

$$\hat{w}_{t+1} = \hat{w}_t - \gamma_t \frac{1}{n} \sum_{i=1}^n x_i \ell'(x_i^\top w - y_i),$$

Then $w_t \in \text{Range}(\hat{X}^\top)$, and if $\ell(0) = 0$ and $n > d$ then $\hat{w}_t \rightarrow \hat{w}^\dagger$ (why?).

Classification

Consider $\ell(x^\top wy)$ and

$$\hat{w}_{t+1} = \hat{w}_t - \gamma_t \frac{1}{n} \sum_{i=1}^n x_i \ell'(x_i^\top wy_i).$$

Then $w_t \in \text{Range}(\hat{X}^\top)$. Convergence to max margin solutions can be ensured in some cases.

Terminology: regularization and pseudosolutions?

- ▶ In signal processing minimal norm solutions are called regularization.
- ▶ In classical regularization theory, they are called pseudosolutions.
- ▶ Regularization refers to a family of solutions converging to pseudosolutions, e.g. Tikhonov.

Outline

The implicit bias of gradient descent

Gradient descent stability and early stopping

Back for more regularization

- ▶ There are problems in which $\hat{w}^\dagger$ is unstable.
- ▶ Regularization: define a family of solutions + select one ensuring stability.

Tikhonov regularization

$$\hat{w}^\lambda \rightarrow \hat{w}^\dagger,$$

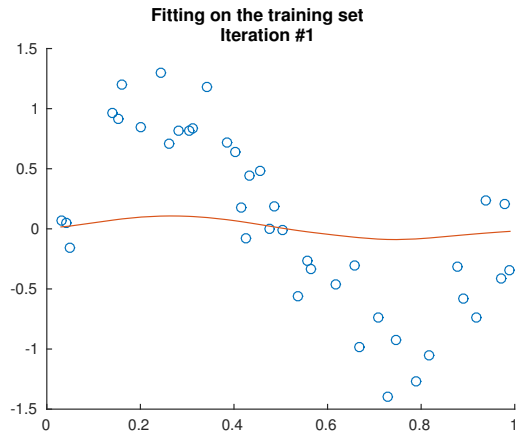
as $\lambda \rightarrow 0$. Choose $\lambda \neq 0$ if needed.

Regularization by gradient descent?

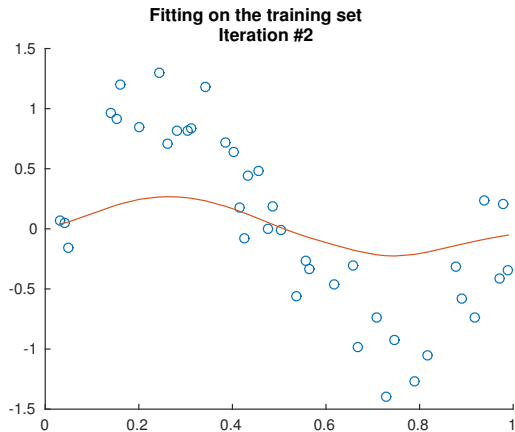
Gradient descent converges to the minimal norm solution, but:

- ▶ does it define a family of meaningful regularized solutions?
- ▶ Where is the regularization parameter?

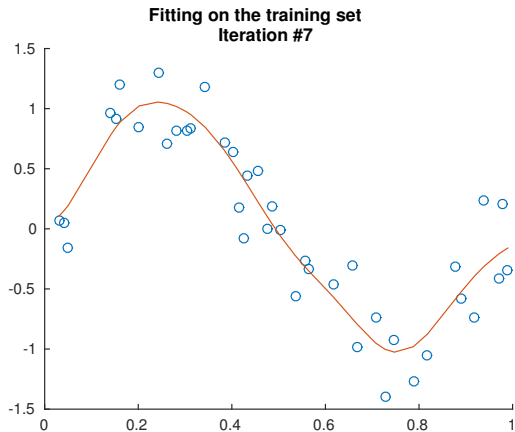
An intuition: early stopping



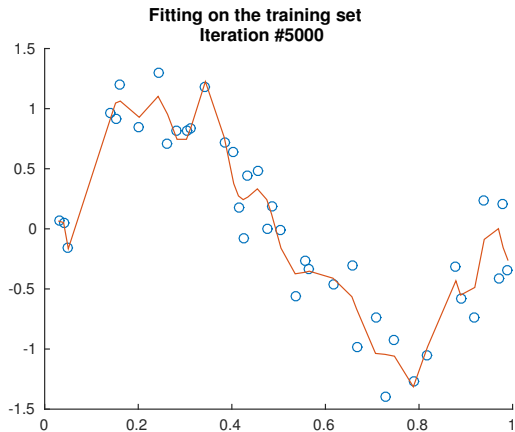
An intuition: early stopping



An intuition: early stopping



An intuition: early stopping



Is there a way to formalize this intuition?

Interlude: geometric series

Recall for $|a| < 1$

$$\sum_{j=0}^{\infty} a^j = (1 - a)^{-1}, \quad \sum_{j=0}^t a^j = (1 - a^{t+1})(1 - a)^{-1}.$$

Interlude: geometric series

Recall for $|a| < 1$

$$\sum_{j=0}^{\infty} a^j = (1 - a)^{-1}, \quad \sum_{j=0}^t a^j = (1 - a^{t+1})(1 - a)^{-1}.$$

Equivalently for $|b| < 1$

$$\sum_{j=0}^{\infty} (1 - b)^j = b^{-1}, \quad \sum_{j=0}^t (1 - b)^j = (1 - (1 - b)^{t+1})b^{-1}.$$

Interlude II: Neumann series

Assume $I - A$ invertible matrix and $\|A\| < 1$

$$\sum_{j=0}^{\infty} A^j = (I - A)^{-1}, \quad \sum_{j=0}^t A^j = (I - A^{t+1})(I - A)^{-1}.$$

²Argument can be extended to pseudoinverses.

Interlude II: Neumann series

Assume $I - A$ invertible matrix and $\|A\| < 1$

$$\sum_{j=0}^{\infty} A^j = (I - A)^{-1}, \quad \sum_{j=0}^t A^j = (I - A^{t+1})(I - A)^{-1}.$$

Equivalently for B invertible² and $\|B\| < 1$

$$\sum_{j=0}^{\infty} (I - B)^j = B^{-1}, \quad \sum_{j=0}^t (I - B)^j = (I - (I - B)^{t+1})B^{-1}.$$

²Argument can be extended to pseudoinverses.

Rewriting GD

For $w_0 = 0$, by induction

$$\hat{w}_{t+1} = \hat{w}_t - \frac{\gamma}{n} \hat{X}^\top (\hat{X} \hat{w}_t - \hat{Y}),$$

can be written as

$$\hat{w}_{t+1} = \frac{\gamma}{n} \sum_{j=0}^t (I - \gamma \hat{X}^\top \hat{X})^j \hat{X}^\top \hat{Y}.$$

Rewriting GD (cont.)

► Write

$$\hat{\mathbf{w}}_{t+1} = \hat{\mathbf{w}}_t - \frac{\gamma}{n} \hat{\mathbf{X}}^\top (\hat{\mathbf{X}} \hat{\mathbf{w}}_t - \hat{\mathbf{Y}}) = \left(\mathbf{I} - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}} \right) \hat{\mathbf{w}}_t + \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{Y}}.$$

Rewriting GD (cont.)

- Write

$$\hat{\mathbf{w}}_{t+1} = \hat{\mathbf{w}}_t - \frac{\gamma}{n} \hat{\mathbf{X}}^\top (\hat{\mathbf{X}} \hat{\mathbf{w}}_t - \hat{\mathbf{Y}}) = (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}}) \hat{\mathbf{w}}_t + \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{Y}}.$$

- Assume

$$\hat{\mathbf{w}}_t = \frac{\gamma}{n} \sum_{j=0}^{t-1} (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}})^j \hat{\mathbf{X}}^\top \hat{\mathbf{Y}}.$$

Rewriting GD (cont.)

- Write

$$\hat{\mathbf{w}}_{t+1} = \hat{\mathbf{w}}_t - \frac{\gamma}{n} \hat{\mathbf{X}}^\top (\hat{\mathbf{X}} \hat{\mathbf{w}}_t - \hat{\mathbf{Y}}) = (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}}) \hat{\mathbf{w}}_t + \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{Y}}.$$

- Assume

$$\hat{\mathbf{w}}_t = \frac{\gamma}{n} \sum_{j=0}^{t-1} (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}})^j \hat{\mathbf{X}}^\top \hat{\mathbf{Y}}.$$

- Then

$$\begin{aligned} \hat{\mathbf{w}}_{t+1} &= (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}}) \frac{\gamma}{n} \sum_{j=0}^{t-1} (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}})^j \hat{\mathbf{X}}^\top \hat{\mathbf{Y}} + \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{Y}} \\ &= \frac{\gamma}{n} \sum_{j=0}^t (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}})^j \hat{\mathbf{X}}^\top \hat{\mathbf{Y}}. \end{aligned}$$

Neumann series and GD

$$\hat{\mathbf{w}}_{t+1} = \frac{\gamma}{n} \sum_{j=0}^t (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}})^j \hat{\mathbf{X}}^\top \hat{\mathbf{y}}.$$

GD is a truncated power series approximation of the pseudoinverse!

If γ is such that³ $\left\| I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}} \right\| < 1$, then for large t

$$\frac{\gamma}{n} \sum_{j=0}^t (I - \frac{\gamma}{n} \hat{\mathbf{X}}^\top \hat{\mathbf{X}})^j \hat{\mathbf{X}}^\top \approx \hat{\mathbf{X}}^\dagger$$

and we recover $\hat{\mathbf{w}}_t \rightarrow \hat{\mathbf{w}}^\dagger$.

³Compare to classic conditions.

Stability properties of GD

For any t

$$\hat{w}_t = (I - (I - \frac{\gamma}{n} \hat{X}^\top \hat{X})^t) (\hat{X}^\top \hat{X})^{-1} \hat{X}^\top \hat{Y}$$

(assume invertibility for simplicity).

Then

$$\underbrace{\hat{w}_t \approx (\hat{X}^\top \hat{X})^{-1} \hat{X}^\top \hat{Y}}_{\text{large } t},$$

$$\underbrace{\hat{w}_t \approx \frac{\gamma}{n} \hat{X}^\top \hat{Y}}_{\text{small } t}.$$

Compare to Tikhonov $\hat{w}_\lambda = (\hat{X}^\top \hat{X} + \lambda n I)^{-1} \hat{X}^\top \hat{Y}$

$$\underbrace{\hat{w}_\lambda \approx (\hat{X}^\top \hat{X})^{-1} \hat{Y}}_{\text{small } \lambda},$$

$$\underbrace{\hat{w}_\lambda \approx \lambda n \hat{X}^\top \hat{Y}}_{\text{large } \lambda}.$$

Spectral view and filtering

Recall for Tikhonov

$$\hat{w}^\lambda = \sum_{j=1}^r \frac{s_j}{s_j^2 + \lambda} (u_j^\top \hat{Y}) v_j.$$

Spectral view and filtering

Recall for Tikhonov

$$\hat{w}^\lambda = \sum_{j=1}^r \frac{s_j}{s_j^2 + \lambda} (u_j^\top \hat{Y}) v_j.$$

For GD

$$\hat{w}^\lambda = \sum_{j=1}^r \frac{(1 - (1 - \frac{\gamma}{n} s_j^2)^t)}{s_j} (u_j^\top \hat{Y}) v_j.$$

Spectral view and filtering

Recall for Tikhonov

$$\hat{w}^\lambda = \sum_{j=1}^r \frac{s_j}{s_j^2 + \lambda} (u_j^\top \hat{Y}) v_j.$$

For GD

$$\hat{w}^\lambda = \sum_{j=1}^r \frac{(1 - (1 - \frac{\gamma}{n} s_j^2)^t)}{s_j} (u_j^\top \hat{Y}) v_j.$$

Both methods can be seen as performing spectral filtering

$$\hat{w}^\lambda = \sum_{j=1}^r F(s_j) (u_j^\top \hat{Y}) v_j,$$

for some suitable filter function F .

Implicit regularization and early stopping

The stability of GD decreases with t , i.e. higher condition number for

$$(I - (I - \frac{\gamma}{n} \hat{X}^\top \hat{X})^t)(\hat{X}^\top \hat{X})^{-1} \hat{X}^\top.$$

Early-stopping the iteration as a (implicit) regularization effect.

Summary so far

$$\hat{\mathbf{w}}_{t+1} = \hat{\mathbf{w}}_t - \frac{\gamma}{n} \hat{\mathbf{X}}^\top (\hat{\mathbf{X}} \hat{\mathbf{w}} - \hat{\mathbf{Y}}) = \frac{\gamma}{n} \sum_{j=0}^t (I - \gamma \hat{\mathbf{X}}^\top \hat{\mathbf{X}})^j \hat{\mathbf{X}}^\top \hat{\mathbf{Y}}.$$

- Implicit bias: gradient descent converges to the minimal norm solution.
- Stability: the number of iteration is a regularization parameter.

Name game: gradient descent, Landweber iteration, L^2 -Boosting.

A bit of history

These ideas are fashionable now but have also a long history.

- ▶ The idea that iterations converge to pseudosolutions is from the 50's.
- ▶ The observation that iterations control stability dates back at least to the 80's.

Classic name is iterative regularization (there are books about it).

Why is it back in fashion?

- ▶ Early stopping is used as a heuristic while training neural nets.
- ▶ Convergence to minimal norm solutions could help understanding generalization in deep learning?
- ▶ New perspective on algorithm design merging statistics and optimization.

Statistics meets optimization

- ▶ Training time= model complexity?
- ▶ Iterations control statistical accuracy *and* numerical complexity.
- ▶ This kind of regularization is also called computational or algorithmic.

Beyond linear least squares

- ▶ Other forms of optimization?
- ▶ Other loss functions?
- ▶ Other class of functions?
- ▶ Other norms?

Other forms of optimization

Largely unexplored, there are some results on:

- ▶ Accelerated methods and conjugate gradient.
- ▶ Stochastic/incremental gradient methods.

Mostly for least squares.

It is clear that other parameters control regularization/stability, e.g step-size, mini-batch-size, averaging etc.

Other loss functions

There are some results.

For ℓ convex, let

$$\hat{L}(w) = \frac{1}{n} \sum_{i=1}^n \ell(y_i, w^\top x_i).$$

The gradient/subgradient descent iteration is

$$\hat{w}_{t+1} = \hat{w}_t - \gamma_t \nabla \hat{L}(\hat{w}_t).$$

Other loss functions (cont.)

$$\hat{w}_{t+1} = \hat{w}_t - \gamma_t \nabla \hat{L}(\hat{w}_t)$$

An intuition: note that, if $\sup_t \left\| \nabla \hat{L}(\hat{w}_t) \right\| \leq B$,

$$\|\hat{w}_t\| \leq \sum_t \gamma_t B,$$

the number of iterations/stepsize control the norm of the iterates.

Representer theorem for GD

$$\hat{w}_{t+1} = \hat{w}_t - \gamma_t \frac{1}{n} \sum_{i=1}^n x_i \ell'(x_i^\top \hat{w}_t, y_i).$$

It is easy to show that by induction that

$$\hat{w}_t = \sum_{i=1}^n x_i (c_t)_i, \Leftrightarrow \hat{f}_t(x) = x^\top \hat{w}_t = \sum_{i=1}^n x^\top x_i (c_t)_i$$

$$c_{t+1} = \hat{w}_t - \gamma_t \frac{1}{n} \sum_{i=1}^n e_i \ell'(x_i^\top \hat{w}_t, y_i),$$

This will give us kernels! Neural nets is trickier...

Other norms

Some results.

- ▶ Gradient descent needs be replaced to bias iterations towards desired norms.
- ▶ Bregman iterations, mirror descent, proximal gradients can be used.

Summing up

- ▶ A different way to design algorithms.
- ▶ Implicit/iterative regularization.
- ▶ Iterative regularization for least squares.
- ▶ Extensions.