

# On Regression: From Least Squares to Kernels and Neural Tangent Kernels

Advanced Machine Learning - ETH Zürich 2026

**Reading guide.** These notes expand the blackboard derivations. We assume undergrad-level linear algebra, probability, and vector calculus. Shapes are shown next to key equations to keep dimensions straight.

## Contents

|          |                                                                            |          |
|----------|----------------------------------------------------------------------------|----------|
| <b>1</b> | <b>Problem setup and notation</b>                                          | <b>1</b> |
| <b>2</b> | <b>Act I: Linear regression is (can be) unstable</b>                       | <b>2</b> |
| 2.1      | MLE equals OLS                                                             | 2        |
| 2.2      | Bias–variance–noise decomposition (test error sketch)                      | 2        |
| 2.3      | Instability via SVD                                                        | 2        |
| <b>3</b> | <b>Act II: Stabilizing with regularization (Bayesian view of ridge)</b>    | <b>2</b> |
| 3.1      | Gaussian prior and posterior                                               | 2        |
| 3.2      | Variance shrinkage in eigen-directions                                     | 3        |
| <b>4</b> | <b>Act III: Beyond linear—features and kernels</b>                         | <b>3</b> |
| 4.1      | Feature maps and kernel trick                                              | 3        |
| 4.2      | A polynomial (countably infinite) feature map yielding an RBF kernel       | 3        |
| 4.3      | Kernelized regression                                                      | 3        |
| <b>5</b> | <b>Act IV: Neural networks &amp; gradient descent as kernel regression</b> | <b>3</b> |
| 5.1      | Model and training                                                         | 3        |
| 5.2      | Linearization/Taylor step at initialization                                | 4        |
| 5.3      | Training dynamics and the NTK                                              | 4        |
| <b>6</b> | <b>Summary and takeaways</b>                                               | <b>4</b> |

## 1 Problem setup and notation

- Data:  $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$  with  $\mathbf{x}_i \in \mathbb{R}^d$ ,  $y_i \in \mathbb{R}$ ; stack  $\mathbf{X} = [\mathbf{x}_1^\top; \dots; \mathbf{x}_n^\top] \in \mathbb{R}^{n \times d}$  and  $\mathbf{y} \in \mathbb{R}^n$ .
- Loss: mean squared error (MSE):  $L(f) = \frac{1}{n} \sum_{i=1}^n (y_i - f(\mathbf{x}_i))^2$ .
- Noise model for sections 2–3:  $y_i = f^*(\mathbf{x}_i) + \varepsilon_i$ ,  $\varepsilon_i \sim \mathcal{N}(0, \sigma^2)$  i.i.d.

## 2 Act I: Linear regression is (can be) unstable

### 2.1 MLE equals OLS

Assume a linear nonparametric model (L-NPM):  $f_{\beta(\mathbf{x})=\mathbf{x}^\top \beta}$  with  $\beta \in \mathbb{R}^d$ . The likelihood is

$$p(\mathbf{y} \mid \beta) \propto \exp\left(-\frac{1}{2\sigma^2} \|\mathbf{y} - \mathbf{X}\beta\|^2\right).$$

Maximizing the likelihood is equivalent to minimizing MSE. The normal equations:

$$\mathbf{X}^\top \mathbf{X} \hat{\beta} = \mathbf{X}^\top \mathbf{y}.$$

**Proposition 1** (OLS/Pseudoinverse form). *If  $\mathbf{X}$  has full column rank,*

$$\hat{\beta} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}.$$

*In general,  $\hat{\beta} = \mathbf{X}^\dagger \mathbf{y}$ , the minimum-norm least-squares solution.*

### 2.2 Bias–variance–noise decomposition (test error sketch)

State the standard decomposition for a new point  $\mathbf{x}$ :

$$\mathbb{E}\left[(y - \mathbf{x}^\top \hat{\beta})^2\right] = \underbrace{\left(f^*(\mathbf{x}) - \mathbb{E}[\mathbf{x}^\top \hat{\beta}]\right)^2}_{\text{bias}^2} + \underbrace{\text{Var}(\mathbf{x}^\top \hat{\beta})}_{\text{variance}} + \underbrace{\sigma^2}_{\text{irreducible noise}}.$$

### 2.3 Instability via SVD

Let the thin SVD of  $\mathbf{X}$  be  $\mathbf{X} = UDV^\top$  with singular values  $\{s_i\}$  on  $D$ . Then

$$\hat{\beta} = VD^{-1}U^\top \mathbf{y} \quad \text{and} \quad \text{Var}(\hat{\beta}) = \sigma^2 (\mathbf{X}^\top \mathbf{X})^{-1} = \sigma^2 VD^{-2}V^\top = \sigma^2 \sum_i \frac{1}{s_i^2} v_i v_i^\top.$$

**Remark 1** (Why OLS can blow up). *In high dimension or collinearity, some  $s_i \approx 0$ , making  $\text{Var}(\hat{\beta})$  huge. This is the core instability motivating regularization.*

## 3 Act II: Stabilizing with regularization (Bayesian view of ridge)

### 3.1 Gaussian prior and posterior

Place a prior  $\beta \sim \mathcal{N}(0, \tau^2 \mathbf{I})$ . The posterior is Gaussian:

$$p(\beta \mid \mathbf{y}) \propto \exp\left(-\frac{1}{2\sigma^2} \|\mathbf{y} - \mathbf{X}\beta\|^2 - \frac{1}{2\tau^2} \|\beta\|^2\right).$$

**Proposition 2** (MAP = Ridge). *With  $\lambda = \sigma^2/\tau^2$ ,*

$$\hat{\beta}_{MAP} = \arg \min_{\beta} \frac{1}{n} \|\mathbf{y} - \mathbf{X}\beta\|^2 + \lambda \|\beta\|^2 = (\mathbf{X}^\top \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^\top \mathbf{y}.$$

### 3.2 Variance shrinkage in eigen-directions

Using the SVD of  $\mathbf{X}$ ,

$$\text{Var}(\mathbf{x}^\top \hat{\boldsymbol{\beta}}_{\text{ridge}}) = \sigma^2 \sum_i \frac{s_i^2}{(s_i^2 + \lambda)^2} (\mathbf{x}^\top \mathbf{v}_i)^2,$$

and for the coefficient vector

$$\text{Var}(\hat{\boldsymbol{\beta}}_{\text{ridge}}) = \sigma^2 \sum_i \frac{s_i^2}{(s_i^2 + \lambda)^2} \mathbf{v}_i \mathbf{v}_i^\top.$$

**Remark 2.** *Small singular directions ( $s_i \approx 0$ ) are heavily shrunk, trading variance for bias and restoring stability.*

## 4 Act III: Beyond linear—features and kernels

### 4.1 Feature maps and kernel trick

Let  $\phi : \mathbb{R}^d \rightarrow \mathcal{H}$  be (possibly infinite) features and  $\Phi \in \mathbb{R}^{n \times p}$  be the design with rows  $-\phi(\mathbf{x}_i)^\top$  (sign chosen to match the board algebra). Linear regression in feature space:

$$\hat{\boldsymbol{\beta}} = \Phi^\top (\Phi \Phi^\top + \lambda \mathbf{I})^{-1} \mathbf{y}, \quad \hat{f}(x_\star) = \phi(x_\star)^\top \hat{\boldsymbol{\beta}} = k(x_\star)^\top (K + \lambda \mathbf{I})^{-1} \mathbf{y},$$

where  $K_{ij} = k(\mathbf{x}_i, \mathbf{x}_j) = \langle \phi(\mathbf{x}_i), \phi(\mathbf{x}_j) \rangle$  and  $k(x_\star) = (k(x_\star, \mathbf{x}_i))_{i=1}^n$ .

### 4.2 A polynomial (countably infinite) feature map yielding an RBF kernel

Define multi-index  $\alpha \in \mathbb{N}^d$ ,  $|\alpha| = \sum_j \alpha_j$ ,  $\mathbf{x}^\alpha = \prod_j x_j^{\alpha_j}$ ,  $\alpha! = \prod_j \alpha_j!$ . Consider the (normalized) map

$$\phi(\mathbf{x}) = e^{-\frac{1}{2}\|\mathbf{x}\|^2} \left( \frac{\mathbf{x}^\alpha}{\sqrt{\alpha!}} \right)_{\alpha \in \mathbb{N}^d}.$$

Then

$$\langle \phi(\mathbf{x}), \phi(\mathbf{x}') \rangle = \exp\left(-\frac{1}{2}\|\mathbf{x}\|^2 - \frac{1}{2}\|\mathbf{x}'\|^2\right) \sum_{\alpha} \frac{(\mathbf{x}^\top \mathbf{x}')^{|\alpha|}}{\alpha!} = \exp\left(-\frac{1}{2}\|\mathbf{x} - \mathbf{x}'\|^2\right),$$

the RBF kernel.<sup>1</sup>

### 4.3 Kernelized regression

State the final predictor (with or without  $\lambda$ ):

$$\boxed{\hat{f}(x_\star) = k(x_\star)^\top (K + \lambda \mathbf{I})^{-1} \mathbf{y}}.$$

Discuss computational trade-offs:  $O(n^3)$  to invert  $K$ , independence from feature dimension.

## 5 Act IV: Neural networks & gradient descent as kernel regression

### 5.1 Model and training

One-hidden-layer network:

$$f_{\boldsymbol{\theta}}(\mathbf{x}) = \frac{1}{\sqrt{m}} \sum_{r=1}^m \alpha_r \varphi_r(\mathbf{x}), \quad \boldsymbol{\theta} = \{\alpha_{r,r}\}_{r=1}^m, \quad \boldsymbol{\theta}_0 \sim \mathcal{N}(0, \omega^2 \mathbf{I}).$$

---

<sup>1</sup>Use the multinomial theorem and the series for  $\exp(t)$ .

Gradient descent on MSE with step  $\eta$ :

$$\boldsymbol{\theta}_{t+1} = \boldsymbol{\theta}_t - \eta \nabla_{\boldsymbol{\theta}} \frac{1}{2n} \|\mathbf{y} - f_t\|^2, \quad \nabla_{\boldsymbol{\theta}} \frac{1}{2n} \|\mathbf{y} - f_t\|^2 = -\frac{1}{n} \Phi_t^\top (\mathbf{y} - f_t),$$

where  $f_t = (f_{\boldsymbol{\theta}_t}(\mathbf{x}_i))_{i=1}^n$  and

$$\Phi_t \in \mathbb{R}^{n \times p}, \quad (\Phi_t)_{i,:} = \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}_t}(\mathbf{x}_i)^\top.$$

## 5.2 Linearization/Taylor step at initialization

First-order Taylor around  $\boldsymbol{\theta}_0$ :

$$f_{\boldsymbol{\theta}}(\mathbf{x}) \approx f_{\boldsymbol{\theta}_0}(\mathbf{x}) + \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}_0}(\mathbf{x})^\top (\boldsymbol{\theta} - \boldsymbol{\theta}_0).$$

Stacked over the training points:

$$f \approx f_0 + \Phi_0 (\boldsymbol{\theta} - \boldsymbol{\theta}_0).$$

Solving the least-squares relation for  $(\boldsymbol{\theta} - \boldsymbol{\theta}_0)$  gives (Moore–Penrose):

$$\boldsymbol{\theta} - \boldsymbol{\theta}_0 \approx \Phi_0^\top (\Phi_0 \Phi_0^\top)^{-1} (f - f_0).$$

For a test point  $x_*$ ,

$$f(x_*) \approx f_0(x_*) + \underbrace{\nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}_0}(x_*)^\top \Phi_0^\top}_{k(x_*)^\top} (\Phi_0 \Phi_0^\top)^{-1} (f - f_0).$$

## 5.3 Training dynamics and the NTK

Define the (empirical) Neural Tangent Kernel (NTK) matrix

$$K := \Phi_0 \Phi_0^\top \in \mathbb{R}^{n \times n}, \quad K_{ij} = \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}_0}(\mathbf{x}_i)^\top \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}_0}(\mathbf{x}_j).$$

Under the linearized model and small  $\eta$ ,

$$f_{t+1} = f_t - \eta K (f_t - \mathbf{y}).$$

Hence

$$f_t = \mathbf{y} + (I - \eta K)^t (f_0 - \mathbf{y}) \xrightarrow{t \rightarrow \infty} \mathbf{y} \quad (\text{on the training set if } \eta < 2/\lambda_{\max}(K)).$$

At a test point,

$$\boxed{\hat{f}(x_*) = k(x_*)^\top K^{-1} \mathbf{y}}, \quad k(x_*)_i = \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}_0}(x_*)^\top \nabla_{\boldsymbol{\theta}} f_{\boldsymbol{\theta}_0}(\mathbf{x}_i).$$

**Remark 3** (Wide-network limit (the “hard claims”)). *As width  $m \rightarrow \infty$ , with standard random initialization,  $K$  and  $k(\cdot)$  concentrate around deterministic limits; gradient descent therefore emulates kernel regression with the NTK.*

## 6 Summary and takeaways

- OLS is MLE for Gaussian noise but is unstable when  $\mathbf{X}$  is ill-conditioned or rank-deficient.
- Ridge (MAP with a Gaussian prior) shrinks high-variance directions and stabilizes estimation.
- Infinite feature maps + the kernel trick yield predictors  $\hat{f}(x_*) = k(x_*)^\top (K + \lambda I)^{-1} \mathbf{y}$  without explicit features.
- In the linearized regime, gradient descent training of wide neural nets performs kernel regression with the NTK.

## Appendix (optional material to expand)

- A. Full derivation of the bias–variance decomposition.
- B. Conjugate Gaussian identities used in Section 3.
- C. Derivation of the RBF kernel from the polynomial feature map (multi-index calculus).
- D. Convergence condition for the linearized GD dynamics ( $\eta < 2/\lambda_{\max}(K)$ ).
- E. Notes on computational complexity and numerics (Cholesky for  $(K + \lambda I)^{-1}$ , conditioning, centering).