

Exercises on Optimization for Machine Learning

M2 IASD 2025-2026

Version 2.0 - October 15, 2025*



1. Basics of optimization

Exercise 1.1: Support vector machines

Given a dataset $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$ with $\mathbf{x}_i \in \mathbb{R}^d$ and $y_i \in \{-1, +1\}$, finding a binary linear classifier for the data using support vector machines amounts to considering the problem

$$\underset{\mathbf{w} \in \mathbb{R}^d}{\text{minimize}} \frac{1}{n} \sum_{i=1}^n \max \{1 - y_i \mathbf{x}_i^T \mathbf{w}, 0\} + \frac{\lambda}{2} \|\mathbf{w}\|_2^2, \quad (1)$$

where $\lambda \geq 0$.

Show that the objective of (1) is a convex function for any $\lambda \geq 0$, and that it is λ -strongly convex when $\lambda > 0$.

Exercise 1.2: Least-squares problems

Let $\mathbf{x} \in \mathbb{R}^d$ such that $\|\mathbf{x}\|_2 \neq 0$ and $\mathbf{y} \in \mathbb{R}^d$. The goal of the exercise is to compute a linear model that best fits $\mathbf{x}$ to $\mathbf{y}$, and to show the interest of using overparameterization in that setting.

a) Consider first the one-dimensional problem

$$\underset{w \in \mathbb{R}}{\text{minimize}} \frac{1}{2} \|w\mathbf{x} - \mathbf{y}\|_2^2. \quad (2)$$

Is problem (2) convex? Is its optimal value equal to 0?

b) Consider now the matrix problem

$$\underset{\mathbf{W} \in \mathbb{R}^{d \times d}}{\text{minimize}} \frac{1}{2} \|\mathbf{W}\mathbf{x} - \mathbf{y}\|_2^2. \quad (3)$$

Reformulate the problem so as to show that the objective function is convex in the sense defined in class, i.e. for vector-valued functions.

*Thanks to the students who spotted typos in earlier versions!

c) Show that

$$\underset{\mathbf{W} \in \mathbb{R}^{d \times d}}{\text{minimize}} \frac{1}{2} \|\mathbf{W}\mathbf{x} - \mathbf{y}\|_2^2 = 0$$

and that the problem does not have a unique global minimum.

Exercise 1.3: Strong convexity

Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be $\mathcal{C}^1$ and μ -strongly convex, and denote by $\mathbf{w}^*$ the minimum of f .

a) For any $\mathbf{w} \in \mathbb{R}^d$, justify that the function

$$\varphi_{\mathbf{w}} : \mathbf{z} \mapsto f(\mathbf{w}) + \nabla f(\mathbf{w})^T(\mathbf{z} - \mathbf{w}) + \frac{\mu}{2} \|\mathbf{z} - \mathbf{w}\|^2$$

is $\mathcal{C}^1$ and strongly convex.

b) Using $\nabla \varphi_{\mathbf{w}}(\mathbf{z}) = \nabla f(\mathbf{w}) + \mu(\mathbf{z} - \mathbf{w})$ for any $\mathbf{z} \in \mathbb{R}^d$, compute $\min_{\mathbf{z}} \varphi_{\mathbf{w}}(\mathbf{z})$ and $\operatorname{argmin}_{\mathbf{z}} \varphi_{\mathbf{w}}(\mathbf{z})$.

c) Conclude from the previous questions that

$$\|\nabla f(\mathbf{w})\|_2^2 \geq 2\mu (f(\mathbf{w}) - f(\mathbf{w}^*)).$$

Exercise 1.4: Co-coercivity

Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$ be $\mathcal{C}_L^{1,1}$ and convex. Suppose that $\mathbf{w}^* \in \operatorname{argmin}_{\mathbf{w}} f(\mathbf{w})$ and let $f^* = f(\mathbf{w}^*)$.

a) Let $\mathbf{w} \in \mathbb{R}^d$. Show that $f(\mathbf{w}) - f(\mathbf{w}^*) \geq \frac{1}{2L} \|\nabla f(\mathbf{w})\|_2^2$.

b) Let $(\mathbf{w}, \mathbf{v}) \in (\mathbb{R}^d)^2$. Show that

$$(\nabla f(\mathbf{v}) - \nabla f(\mathbf{w}))^T(\mathbf{v} - \mathbf{w}) \geq \frac{1}{L} \|\nabla f(\mathbf{v}) - \nabla f(\mathbf{w})\|_2^2.$$

Consider $\mathbf{z} \mapsto f(\mathbf{z}) - \nabla f(\mathbf{v})^T \mathbf{z}$ and $\mathbf{z} \mapsto f(\mathbf{z}) - \nabla f(\mathbf{w})^T \mathbf{z}$.

2. Derivatives and differentiation

Exercise 2.1: A simple function (Adapted from 2019-2020 exam)

Consider the one-dimensional function $f : \mathbb{R} \rightarrow \mathbb{R}$ defined by

$$f(w) = \sqrt{\sin(w)} + \sin(w). \quad (4)$$

a) Write a directed acyclic graph representing the computation of $f(w)$.

b) Write down the instructions for computing $f'(w)$ in forward mode. Use it to evaluate $f'(\frac{\pi}{2})$.

c) Write down the instructions for computing $f'(w)$ in backward mode. Use it to evaluate $f'(\frac{\pi}{2})$.

Exercise 2.2: Sigmoid-type functions (Adapted from 2020-2021 exam)

Given $\mathbf{x} \in \mathbb{R}^d$ and $y \in [0, 1]$, we consider the two functions $f, g : \mathbb{R}^d \rightarrow \mathbb{R}$ defined by

$$f(\mathbf{w}) = y \ln(\sigma(\mathbf{x}^T \mathbf{w})) \quad \text{and} \quad g(\mathbf{w}) = (1 - y) \ln(1 - \sigma(\mathbf{x}^T \mathbf{w})), \quad (5)$$

where $\sigma(t) = \frac{1}{1+e^{-t}}$. Recall that $\sigma'(t) = \sigma(t)(1 - \sigma(t))$.

- Write a computational graph representing the calculation of $f = f(\mathbf{w})$ by considering dot product, $\sigma(\cdot)$, logarithm and product as elementary operations.
- Write down the chain rule for computing $\frac{\partial f}{\partial \mathbf{w}}$ in forward mode.
- Write down the chain rule for computing $\frac{\partial f}{\partial \mathbf{w}}$ in backward mode.
- Write a computational graph representing the calculation of $h = f(\mathbf{w}) + g(\mathbf{w})$ based on the graph from question a).
- Write down the chain rule for computing $\frac{\partial h}{\partial \mathbf{w}}$ in backward mode.

3. Gradient descent

Exercise 3.1: A strongly convex problem

Consider the minimization problem

$$\underset{\mathbf{w} \in \mathbb{R}^d}{\text{minimize}} \quad f(\mathbf{w}) := \frac{1}{2} \|\mathbf{w}\|^2. \quad (6)$$

For any $\mathbf{w} \in \mathbb{R}^d$, we have $\nabla f(\mathbf{w}) = \mathbf{w}$ and $\nabla^2 f(\mathbf{w}) = \mathbf{I}_d$.

- Justify that $f \in \mathcal{C}_1^{1,1}(\mathbb{R}^d)$ and that f is 1-strongly convex.
- Adapt the general convergence rates of gradient descent and accelerated gradient on strongly convex problems to this particular case. What do the results suggest about the problem?
- Write down the first iteration of gradient descent and accelerated gradient for the problem at hand using a stepsize of 1: is the result in agreement with the convergence rates?

Exercise 3.2: Convergence rates

Let $f : \mathbb{R}^d \rightarrow \mathbb{R}$. Suppose that $f \in \mathcal{C}_L^{1,1}(\mathbb{R}^d)$ and that f is μ -strongly convex. Suppose that we apply gradient descent with a constant stepsize $\alpha_k = \frac{1}{\mu+L}$ to this problem.

- Using the descent property stated in class, justify that one has

$$f(\mathbf{w}_{k+1}) \leq f(\mathbf{w}_k) - \frac{2\mu + L}{2(\mu + L)^2} \|\nabla f(\mathbf{w}_k)\|_2^2 \quad (7)$$

holds at every iteration of gradient descent.

- b) Using the previous question, show that gradient descent achieves a linear convergence rate with this stepsize. Are additional assumptions on μ and/or L needed?
- c) Gradient descent with stepsize $\frac{1}{L}$ achieves a linear rate of convergence in $(1 - \frac{\mu}{L})^K$ on this problem. Compare this rate with that obtained in question b).
- d) Accelerated gradient (with the appropriate parameterization for strongly convex functions) achieves a linear rate of convergence in $(1 - \sqrt{\frac{\mu}{L}})^K$ on this problem. Compare this rate with that obtained in question b).

Exercise 3.3: Saddle points

Consider the two-dimensional function $f : \mathbb{R}^2 \rightarrow \mathbb{R}$ defined by

$$f(\mathbf{w}) := w_1^2 (w_1^2 - 2) + w_2^4.$$

- a) For every $\mathbf{w} \in \mathbb{R}^2$, we have $\nabla f(\mathbf{w}) = \begin{bmatrix} 4w_1^3 - 4w_1 \\ 4w_2^3 \end{bmatrix}$. Find all first-order stationary points of f .
- b) The formula for the Hessian matrix gives

$$\forall \mathbf{w} \in \mathbb{R}^d, \nabla^2 f(\mathbf{w}) = \begin{bmatrix} 12w_1^2 - 4 & 0 \\ 0 & 12w_2^2 \end{bmatrix}.$$

Evaluate the Hessian at the points found in the previous question. What can you conclude about the nature of those points?

- c) Without using derivatives, show that the critical points are either global minima or strict saddle points.

Exercise 3.4: Gradient descent with line search

Consider the minimization of a nonconvex, $\mathcal{C}^1$ function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, assuming that f is bounded below by $f_{\text{low}} \in \mathbb{R}$. Suppose that we use gradient descent with a stepsize chosen through backtracking line search. More precisely, given $\theta \in (0, 1)$ and $c \in (0, 1/2)$, we select the stepsize at iteration k as the largest value in $\{\theta^j\}_{j \in \mathbb{N}}$ such that

$$f(\mathbf{w}_k - \theta^j \nabla f(\mathbf{w}_k)) < f(\mathbf{w}_k) - c\theta^j \|\nabla f(\mathbf{w}_k)\|_2^2. \quad (8)$$

- a) Suppose that $\|\nabla f(\mathbf{w}_k)\|_2 > 0$ at iteration k . Show then that the line search terminates within a fixed number of iterations (independent of $\mathbf{w}_k$) and give a bound on the value of the stepsize α_k .
- b) Show that the convergence rate of gradient descent in this setting is $\mathcal{O}(\frac{1}{\sqrt{K}})$, as in the fixed-stepsize case seen in class. Give the precise values for the constants.
- c) What is the maximum number of backtracking steps at each iteration? How does this number illustrate the additional cost of line-search techniques?

Exercise 3.5: Hölder continuity (Adapted from 2021-2022 exam)

In this exercise, we consider a general, unconstrained optimization problem of the form

$$\underset{\mathbf{w} \in \mathbb{R}^d}{\text{minimize}} f(\mathbf{w}), \quad (9)$$

where $f \in \mathcal{C}^1$. Rather than assuming that the gradient is Lipschitz continuous (as in the lectures), we will assume a more general property called *Hölder continuity*. More precisely, for any $\nu \in (0, 1]$ and $L > 0$, we write $f \in \mathcal{C}_L^{1,\nu}$ if $f \in \mathcal{C}^1$ and

$$\forall (\mathbf{v}, \mathbf{w}) \in (\mathbb{R}^d)^2, \quad \|\nabla f(\mathbf{v}) - \nabla f(\mathbf{w})\| \leq L \|\mathbf{v} - \mathbf{w}\|^\nu. \quad (10)$$

We say that ∇f is (L, ν) -Hölder continuous.

For any $f \in \mathcal{C}_L^{1,\nu}$, we can show that

$$\forall (\mathbf{v}, \mathbf{w}) \in (\mathbb{R}^d)^2, \quad f(\mathbf{v}) \leq f(\mathbf{w}) + \nabla f(\mathbf{w})^T (\mathbf{v} - \mathbf{w}) + \frac{L}{1+\nu} \|\mathbf{v} - \mathbf{w}\|_2^2. \quad (11)$$

Moreover, since $f \in \mathcal{C}^1$, we can apply gradient descent to problem (9): the k th iteration of this method is

$$\mathbf{w}_{k+1} = \mathbf{w}_k - \alpha_k \nabla f(\mathbf{w}_k). \quad (12)$$

Part 1 We first study the case $\nu = 1$ (i.e. $f \in \mathcal{C}_L^{1,1}$), corresponding to the gradient being L -Lipschitz continuous.

- Recalling that $f \in \mathcal{C}_L^{1,1}$ for this question, justify the choice $\alpha_k = \frac{1}{L}$ for every k .
- Give the convergence rate result for gradient descent on problem (9), assuming that the function f is nonconvex. What quantity does this rate apply to?
- Give the convergence rate result for gradient descent on problem (9), assuming that the function f is convex. What quantity does this rate apply to?
- Assuming that f is convex, name a method with a better convergence rate than gradient descent, and give the associated rate.

Part 2 We now consider the more general setting $\nu \in (0, 1]$ (i.e. $f \in \mathcal{C}_L^{1,\nu}$) and f nonconvex.

- Based on (11) and your answer to Part 1, justify the choice

$$\alpha_k = \left[\frac{\|\nabla f(\mathbf{w}_k)\|^{1-\nu}}{L} \right]^{\frac{1}{\nu}} \quad (13)$$

for the stepsize used in the iteration (12).

- Using the sequence $\{\alpha_k\}_k$ from the previous question, it is possible to show a convergence rate in $\mathcal{O}\left(\frac{1}{K^{\frac{\nu}{1+\nu}}}\right)$ for gradient descent applied to $f \in \mathcal{C}_L^{1,\nu}$ nonconvex and bounded below. Compare this bound with the results seen in class for the case $\nu = 1$.
- In addition to $f \in \mathcal{C}_L^{1,\nu}$, suppose that f is twice continuously differentiable.
 - Without any convexity assumption on f , are we guaranteed to converge towards a local minimum?
 - What guarantee seen in class can we provide about the limit points of gradient descent?

4. Nonsmooth and regularized optimization

Exercise 4.1: Robust linear regression (Adapted from 2024-2025 exam)

We consider a linear regression problem of the form

$$\underset{\mathbf{w} \in \mathbb{R}^d}{\text{minimize}} \quad f(\mathbf{w}) := \frac{1}{n} \|\mathbf{X}\mathbf{w} - \mathbf{y}\|_1 = \frac{1}{n} \sum_{i=1}^n |\mathbf{x}_i^T \mathbf{w} - y_i|, \quad (14)$$

where $\mathbf{X} = [\mathbf{x}_1 \ \cdots \ \mathbf{x}_n]^T \in \mathbb{R}^{n \times d}$ and $\mathbf{y} \in \mathbb{R}^n$ form the problem data.

- a) Justify that the objective function of (14) is convex.
- b) Let $\mathbf{w}^*$ and $\bar{\mathbf{w}}$ be two solutions of (14). Justify that the vector $\frac{\mathbf{w}^* + \bar{\mathbf{w}}}{2}$ is also a solution of the problem.
- c) Due to the use of the ℓ_1 norm, the function f is not necessarily differentiable.
 - i) Give a condition on the data $(\mathbf{X}, \mathbf{y})$ under which f is actually continuously differentiable.
 - ii) Assuming that the condition from the previous question does not hold, give a point in $\mathbb{R}^d$ at which the function is *not* differentiable.
- d) Given any real $t \in \mathbb{R}$, we let $\text{sgn}(t) = 1$ if $t > 0$, $\text{sgn}(t) = 0$ if $t = 0$ and $\text{sgn}(t) = -1$ if $t < 0$.
 - i) For any $\mathbf{w} \in \mathbb{R}^d$ and $i \in \{1, \dots, n\}$, show that the vector $\text{sgn}(\mathbf{x}_i^T \mathbf{w} - y_i) \mathbf{x}_i$ is a subgradient for the function $\mathbf{w} \mapsto \mathbf{x}_i^T \mathbf{w} - y_i$ at $\mathbf{w}$.
 - ii) Consequently, justify that the vector

$$\mathbf{g}(\mathbf{w}) = \frac{1}{n} \mathbf{X}^T \text{sgn}(\mathbf{X}\mathbf{w} - \mathbf{y}), \quad \text{where} \quad \text{sgn}(\mathbf{X}\mathbf{w} - \mathbf{y}) = \begin{bmatrix} \text{sgn}(\mathbf{x}_1^T \mathbf{w} - y_1) \\ \vdots \\ \text{sgn}(\mathbf{x}_n^T \mathbf{w} - y_n) \end{bmatrix} \quad (15)$$

is a subgradient for f at $\mathbf{w}$.

- iii) Under the assumptions of question 3b), let $\mathbf{w}_0$ be the point considered in that question. Give another example of a subgradient for f at $\mathbf{w}_0$ other than $\mathbf{g}(\mathbf{w}_0)$.
- e) We now consider solving (14) using a subgradient method.
 - i) Write down the iteration of a subgradient method using a constant stepsize and the vector $\mathbf{g}(\mathbf{x})$ defined by (15).
 - ii) Do you expect the function values corresponding to the method's iterates to be monotonically decreasing? Justify your answer.
 - iii) What kind of guarantees seen in class can we provide about that method for a sufficiently small stepsize?

Exercise 4.2: Proximal operator (Adapted from 2022-2023 exam)

In this exercise, we revisit the proximal operator and the proximal gradient method on a specific problem. Given a vector $\mathbf{w} \in \mathbb{R}^d$, we consider

$$\underset{\mathbf{w} \in \mathbb{R}^d}{\text{minimize}} \quad \|\mathbf{w}\|_1 + \frac{1}{2\alpha} \|\mathbf{w} - \mathbf{z}\|_2^2, \quad (16)$$

where $\|\mathbf{w}\|_1 = \sum_{i=1}^d |[\mathbf{w}]_i|$, $\|\mathbf{w}\|_2^2 = \sum_{i=1}^d [\mathbf{w}]_i^2$ and $\alpha > 0$.

- Explain why the objective function of (16) cannot be optimized by gradient-type techniques.
- Justify that problem (16) and

$$\underset{\mathbf{w} \in \mathbb{R}^d}{\text{minimize}} \quad \alpha \|\mathbf{w}\|_1 + \frac{1}{2} \|\mathbf{w} - \mathbf{z}\|_2^2 \quad (17)$$

have the same solution set (argmin). Since both functions are strongly convex, what can be said about this solution set?

- Write down an optimality condition for problem (16).
- In this question, we view problem (16) as computing a proximal operator.
 - Using the definition of the proximal operator, write down the solution of problem (16) as the value of a proximal operator of a certain function.
 - By repeatedly solving instances of problem (16) using the last solution found as $\mathbf{z}$, what algorithm do we obtain?
- In this question, we change perspective and view problem (17) in a composite form, where $\frac{1}{2} \|\mathbf{w} - \mathbf{z}\|_2^2$ is a data-fitting term and $\alpha \|\mathbf{w}\|_1$ is a regularization term.
 - What is the purpose of such a regularization term? Why is it computationally worth using?
 - Write down an iteration of proximal gradient applied to this problem with $\mathbf{w}_k = \mathbf{z}$ and stepsize α . What do you observe then? Is it to be expected?

Exercise 4.3: Regularization and sparsity (Adapted from 2024-2025 exam)

In this exercise, we consider an optimization problem with elastic net regularization of the form

$$\underset{\mathbf{w} \in \mathbb{R}^d}{\text{minimize}} \quad f(\mathbf{w}) + \lambda \left(\|\mathbf{w}\|_1 + \frac{\mu}{2} \|\mathbf{w}\|_2^2 \right), \quad (18)$$

where $\|\cdot\|_2$ denotes the ℓ_2 norm, $\lambda > 0$ and $\mu > 0$.

- Recall the purpose of the regularization terms $\|\cdot\|_1$ and $\|\cdot\|_2^2$ in optimization.
- In this section, we consider the function $\Omega : \mathbf{w} \mapsto \lambda \|\mathbf{w}\|_1 + \frac{\lambda\mu}{2} \|\mathbf{w}\|_2^2$.
 - Recall the definition of the proximal operator associated with Ω , denoted by $\text{prox}_\Omega(\cdot)$. Justify that this particular “prox” can be considered as a function from $\mathbb{R}^d$ to $\mathbb{R}^d$.

ii) Using properties of the proximal operator, one obtains

$$\forall \mathbf{w} \in \mathbb{R}^d, \quad \text{prox}_{\Omega}(\mathbf{w}) = \frac{1}{\lambda\mu + 1} \text{prox}_{\lambda\|\cdot\|_1}(\mathbf{w}). \quad (19)$$

where $\text{prox}_{\lambda\|\cdot\|_1}(\cdot)$ corresponds to the proximal operator for (a multiple of) the ℓ_1 norm. Justify then that $\text{prox}_{\Omega}(\mathbf{w})$ can be computed in closed form.

c) Suppose for this question that f is a continuously differentiable function.

- i) Write down an iteration of the proximal gradient algorithm applied to problem (16) using the prox_{Ω} operator.
- ii) Using the answer from the previous question, justify that the use of elastic net regularization is likely to produce sparse iterates that decay in norm for large values of the regularization parameter λ .