

Partie 1 : calcul du gradient :

Annale 2023-2024:

$$J_1(\vec{w}) = w_0 + x_0 + w_1 x_1 + w_2^2 x_2 + w_3 w_4 x_3^2 - w_4 x_4$$

$\nabla_{\vec{w}} J_1(\vec{w})$ en fonction de $\vec{w}$ et des constantes $(x_0, x_1, x_2, x_3, x_4)$:

def gradient-de- $J_1(w, x)$:

return np.array([1, x[1], 2*w[2]*x[2], w[4]*x[3]**2, w[3]*x[3]**2 - x[4]])

*rq: normalement w est de dimension d } , donc x@d

x est de dimension (n, d)

D = X.shape[1]

N = X.shape[0]

Annale 2024 - 2025:

$$J_1(w) = 1 + x_0 + w_0^2 + x_1^2 w_1 - 2x_2 w_2^2 + x_3^2 w_3 w_4 + x_4^4$$

def gradient-de- $J_1(w, x)$:

return (2*w[0], x[1]**2, -4*x[2]*w[2], x[3]**2*w[4], x[3]**2*w[4], x[3]**2*w[4], x[3]**2*w[4])

np.round(gradient-de- $J_1(w, x)$, 3)

def gradient-de- $J_2(w, x)$:

return x # où $J_2(\vec{w}) = \vec{w} \cdot \vec{x}$

$$l(w, x, y) = \frac{1}{2} (w \cdot x - y)^2$$

$\frac{\partial}{\partial w} l = (w \cdot x - y) \cdot x$

$\downarrow \downarrow \downarrow$
D (N,D) $\uparrow$ * de la dérivée

def gradient-de- $l(w, x, y)$:

return (x @ w - y) * x

Cours / TDs:

I. Qq. notions de base

→ Types d'apprentissage

- supervisé : on connaît la 'vraie' sortie y_n pour chaque entrée x_n (ex : regression, classification)
- non supervisé : pas de sortie y (on cherche une structure dans $P(X)$)

→ représentation des données:

données: $X := \{ \vec{x}^{(1)}, \dots, \vec{x}^{(N)} \} \in \mathbb{R}^{N \times D}$

N : nombre d'échantillons (ligne de X)

D : nombre de features (colonnes de X)

$x_{n,d}$: valeur en ligne n, colonne d

sorties: $Y = (y_1, \dots, y_N)$

GT = valeur réelle de y_n

1) Regression linéaire:

$$f_w(\vec{x}) = w_0 + \sum_{d=1}^D w_{1,d} x_d = w_0 + \vec{w}_1 \cdot \vec{x}$$

prédiction: $\hat{y} = f_w(\vec{x})$

pour un même x , y varie: bruit ou manque d'information

→ on modélise $y = f_w(x) + \epsilon$ ou $\epsilon \sim \mathcal{N}(0, \sigma^2)$
+ $1/\sqrt{D}$

2) Fonction coût et optimisation:

$$J(w, X, Y) = \frac{1}{N} \sum_{n=1}^N (f_w(x_n) - y_n)^2$$

fit = entraîner le modèle = on cherche les param qui minimisent l'erreur de prédiction.

Ex 2 du TD

$$1) J_1(\vec{\theta}) = \frac{1}{2} \|\vec{\theta}\|^2 = \frac{1}{2} \vec{\theta} \cdot \vec{\theta} = \frac{1}{2} (\theta_1^2 + \dots + \theta_D^2)$$

$$\frac{\partial}{\partial \theta_i} J_1 = \theta_i$$

$$\Rightarrow \vec{\nabla}_{\vec{\theta}} J_1 = \vec{\theta}$$

$$2) J_2(\vec{\theta}) = \frac{1}{4} \|\vec{\theta}\|_4^4 = \frac{1}{4} (\theta_1^4 + \dots + \theta_D^4)$$

$$\frac{\partial J_2}{\partial \theta_i} = \theta_i^3$$

$$\Rightarrow \vec{\nabla}_{\vec{\theta}} J_2(\vec{\theta}) = (\theta_1^3, \dots, \theta_D^3)$$

$$3) J_3(\vec{\theta}) = \frac{1}{2} \cdot \frac{1}{N} \sum_{n=1}^N (\vec{\theta} \cdot \vec{x}_n - y_n)^2$$

$$= \frac{1}{2N} \sum_{n=1}^N \left(\sum_{j=1}^D \theta_j x_{nj} - y_n \right)^2$$

$\uparrow$ fonction de coût pour la regression linéaire

$$\rightarrow \sum_{j=1}^D \frac{\partial}{\partial \theta_j} (\theta_j x_{nj} - y_n) = x_{nj}$$

$$\vec{\nabla}_{\vec{\theta}} J_3 = \frac{1}{N} \sum_{n=1}^N (\vec{\theta} \cdot \vec{x}_n - y_n) \vec{x}_n$$

$X \in M_{N,D}(\mathbb{R})$ $Y = \begin{pmatrix} y_1 \\ \vdots \\ y_N \end{pmatrix}$

$$= \begin{pmatrix} x_{1,1} & \dots & x_{1,D} \\ \vdots & & \vdots \\ x_{N,1} & \dots & x_{N,D} \end{pmatrix}$$

$$4) J_4(\vec{\theta}) = \frac{1}{4} \cdot \frac{1}{N} \sum_{n=1}^N (\vec{\theta} \cdot \vec{x}_n - y_n)^4$$

$$= \frac{1}{4} \cdot \frac{1}{N} \sum_{n=1}^N \frac{\partial}{\partial \theta_i} \left[\left(\sum_{j=1}^D \theta_j x_{nj} - y_n \right)^2 \right]$$

$$\vec{\nabla}_{\vec{\theta}} J_4 = \frac{1}{4N} \sum_{n=1}^N 4 (\vec{\theta} \cdot \vec{x}_n - y_n)^3 \vec{x}_n$$

$$= \frac{1}{N} \sum_{n=1}^N (\vec{\theta} \cdot \vec{x}_n - y_n)^3 \vec{x}_n$$

optionnel: $J_5(\vec{\theta}) = \|\vec{\theta}\| = |\theta_1| + \dots + |\theta_D|$

$$\frac{\partial J_5}{\partial \theta_i} = \begin{cases} 1 & \text{si } \theta_i > 0 \\ -1 & \text{si } \theta_i < 0 \end{cases} = \text{sign}(\theta_i)$$

$$\Rightarrow \vec{\nabla}_{\vec{\theta}} J_5 = \text{sign}(\vec{\theta})$$

Partie II : descente de gradient

Annale 2023-2024:

$$f(\vec{w})(\vec{x}_n) = \vec{w} \cdot \vec{x}_n$$

def fw(w, x):

return x @ w

→ regression car les valeurs prises par y sont continues et non discrètes

→ modèle avec biais (caché dans la colonne 1 de X) → c'est un modèle affine

(si 1^{re} colonne = 1 de pour tous les échantillons, alors c'est un modèle affine ($f(x) = \vec{w}'x$ avec $\vec{w}' = (b, w_1, w_2, \dots)$)

calcule le terme $(w \cdot x_n - y_n)^3 x_n$ pour n fixé

def calcul_intermediaire(X, y, n, w):

$$a = (w @ X[n] - y[n]) ** 3 * X[n]$$

(ou: $N = X.shape[0]$

$$a = ((X @ w - y) ** 3).reshape(N, 1) * X[n]$$

→ résultat de shape (N, 1) à la fin (m shape que w toujours!)

→ Calcul du gradient: ici, on code $\nabla_w J(w, X, y) = -a w + \frac{1}{N} \sum_n (w \cdot x_n - y_n)^3 x_n$

def gradient_de_J(w, X, y, alpha):

N = X.shape[0]

D = X.shape[1]

g = np.zeros(D)

for n in range(N):

g += calcul_intermediaire(X, y, n, w)

g /= N

g -= alpha * w

return g

N = X.shape[0]

g = (((X @ w - y) ** 3).reshape(N, 1) * X).sum(axis=0) / N - alpha * w

return g

g = np.mean(((X @ w - y) ** 3).reshape(N, 1) * X, axis=0)

g -= alpha * w

application descente du gradient: def fit(X, y, w0, eta, ITERMAX, alpha):

w = w0.copy() # ici, c'était déjà initialisé

for i in range(ITERMAX):

w -= eta * gradient_de_J(w, X, y, alpha)

return w

II. Partage des données (train / validation split)

X_train, X_valid, y_train, y_valid = sklearn.model_selection.train_test_split(X, y, test_size=0.2, random_state=6)

fonction pour Mean Absolute error $MAE = \frac{1}{N} \sum_{n=1}^N |\vec{w} \cdot \vec{x}_n - y_n|$

def fonction_MAE(w, X, y):

MAE = np.abs(X @ w - y)

MAE = np.mean(MAE)

return MAE

intégration dans la fonction fit

(ajout des param: X_valid, y_valid)

def fit_amelioree(X, y, X_valid, y_valid, w0, eta, ITERMAX, alpha):

MAE_train = np.zeros(ITERMAX)

MAE_valid = np.zeros(ITERMAX)

w = w0.copy()

for i in range(ITERMAX):

w -= eta * gradient_de_J(w, X, y, alpha)

MAE_train[i] = fonction_MAE(w, X, y)

MAE_valid[i] = fonction_MAE(w, X_valid, y_valid)

return w, MAE_train, MAE_valid

* Analyse graphes (notes personnelles)

η trop petit - descente lente

η trop grand - divergence

η bien choisi - convergence stable

ITERMAX trop petit - arrêt prémature

η légèrement trop grand - oscillations

overfitting: train $\downarrow$ valid $\uparrow$

en général:

fonction fw

calcul du gradient

descente de gradient

entraînement avec le

modèle créé.

Annale 2022-2023: # pour $\vec{\nabla}_{\vec{w}} J(\vec{w}, X, Y) = \frac{1}{N} \sum_n (f_w(x_n) - y_n)^3 \vec{x}_n$

def fw(w, x):

return x@w

def gradient_de_J(w, X, Y):

N = X.shape[0]

return $(\frac{1}{N}) * (((fw(w, X) - Y)) ** 3) @ X$

ou: $grad = (((fw(w, X) - Y) ** 3).reshape(N, 1) * X).mean(axis=0)$

initialisation de w avec la loi normale:

def initialiser_w(D):

return np.random.normal(0, D ** -0.5, D)

fonction fit

def fit(X, Y, eta, ITERMAX):

D = X.shape[1]

w = initialiser_w(D)

for i in range(ITERMAX):

w += eta * gradient_de_J(w, X, Y)

return w

fit, avec les valeurs de J et $\eta \|\vec{\nabla}_{\vec{w}} J(\vec{w})\|$

$$J(\vec{w}, X, Y) = \frac{1}{4N} \sum_n (f_w(x_n) - y_n)^4$$

def fit2(X, Y, eta, ITERMAX):

D = X.shape[1]

w = initialiser_w(D)

valeurs_de_J = np.zeros(ITERMAX)

deplacements_de_J = np.zeros(ITERMAX)

for i in range(ITERMAX):

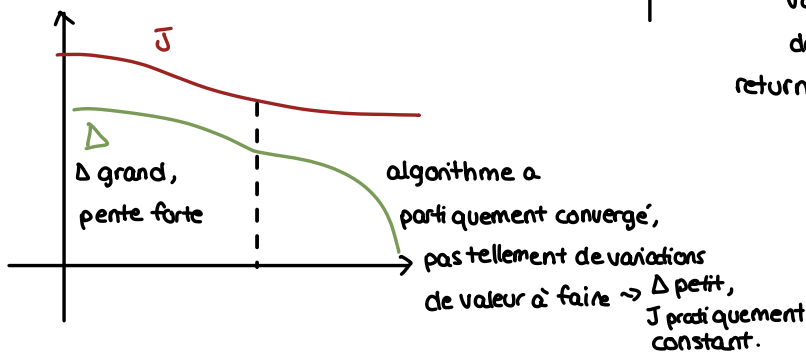
delta = eta * gradient_de_J(w, X, Y)

w += delta

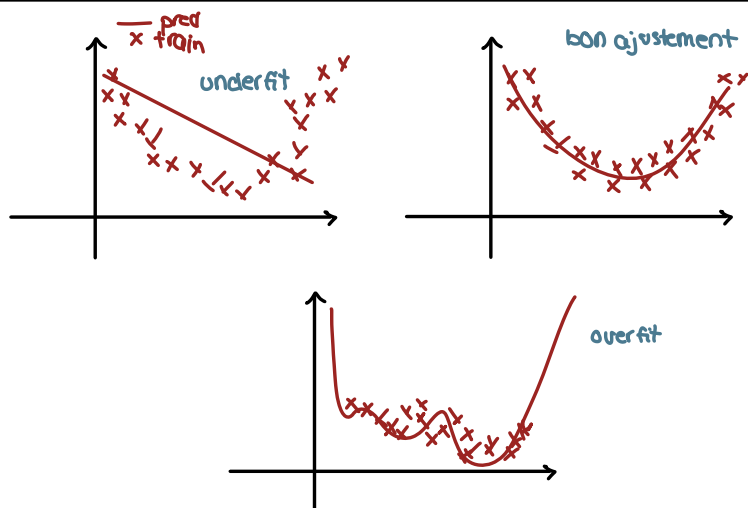
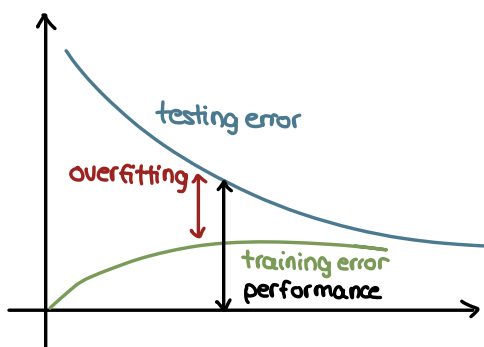
valeurs_de_J[i] = valeur de J(w, X, Y)

deplacements_de_J[i] = np.linalg.norm(delta)

return w, valeurs_de_J, deplacements_de_J



Quelques graphes:



(+ perceptron)

TP: def perceptron FullBatch - version -minimale (X, Y, eta, maxITER=20, plot=None, verbose=None, Loss=None):

w = w0.copy()

for itera in range(maxIter):

ListeDesMalClasses = ((X@w) * Y <= 0)

if ListeDesMalClasses.sum() == 0:

return w

w = w + eta * 1/N * X[ListeDesMalClasses].T @ Y[ListeDesMalClasses]

return w

I: l'ensemble des indices des exemples mal classés

$$J(\vec{w}, X, Y) = \frac{1}{N} \sum_{n \in I} \max(0, -(\vec{w} \cdot \vec{x}^{(n)}) y_n) = \frac{1}{N} \sum_{n \in I} (\vec{w} \cdot \vec{x}^{(n)})$$

~> Mise à jour des paramètres par descente de gradient

η taux d'apprentissage

$$\vec{w} \leftarrow \vec{w} - \eta \vec{\nabla}_{\vec{w}} J(\vec{w}, X, Y)$$

Partie III : optimisation de 1 hyper-paramètre

Annale 2024-2025.

chargement des données (donné)

$X, y = \text{sklearn.datasets.load_digits}(n_class=10, \text{return_X_y}=True)$

$X_trainval, x_test, y_trainval, y_test = \text{sklearn.model_selection.train_test_split}(X, y, \text{test_size}=0.2, \text{random_state}=4)$

plage de valeurs du paramètre

plage-de-valeurs-pour-alpha = $\text{np.logspace}(-2, 4, \text{num}=20)$ [entre -2 et 4]

monModele = $\text{sklearn.linear_model.Ridge}(\text{alpha}=10)$

cross-validation:

nombreDePlisDeCV = 5

$\text{sklearn.model_selection.cross_validate}(\text{monModele}, x_trainval, y_trainval,$
 $\text{cv}=\text{nombreDePlisDeCV}, \text{return_train_score}=True)$

→ renvoi un dictionnaire contenant 4 tableaux :

fit-time → temps d'entraînement par pli

score-time → temps d'évaluation

train-score → score sur l'ensemble d'apprentissage du pli

test-score → score de validation

Chaque tableau contient une valeur par pli.

train-score et test-score pour chaque α dans la plage et les 5 plis

nbplisCV = 5

$\text{trainScore} = \text{np.zeros}((\text{plage-de-valeurs-pour-alpha}.shape[0], \text{nbdeplisCV}))$

$\text{validScore} = \text{np.zeros}((\text{plage-de-valeurs-pour-alpha}.shape[0], \text{nbdeplisCV}))$

for i, α in enumerate(plage-de-valeurs-pour-alpha):

monModele = $\text{sklearn.linear_model.Ridge}(\text{alpha}=\alpha)$

$\text{myCV} = \text{sklearn.model_selection.cross_validate}(\text{monModele}, x_trainval, y_trainval, \text{cv}=\text{nbdeplisCV}, \text{return_train_score}=True)$

$\text{trainScore}[i] = \text{myCV}['\text{train_score}']$

$\text{validScore}[i] = \text{myCV}['\text{test_score}']$

plot des barres d'erreurs

$\text{plt.figure}()$

$\text{plt.errorbar}(\text{plage-de-valeurs-pour-alpha}, \text{np.mean}(\text{trainScore}, \text{axis}=1), \text{label}='train \text{ score}', \text{yerr}=\text{np.std}(\text{trainScore}, \text{axis}=1))$

$\text{plt.errorbar}(\text{plage-de-valeurs-pour-alpha}, \text{np.mean}(\text{validScore}, \text{axis}=1), \text{label}='valid \text{ score}', \text{yerr}=\text{np.std}(\text{validScore}, \text{axis}=1))$

$\text{plt.xlabel}('alpha')$ $\text{plt.legend}()$

$\text{plt.ylabel}('scores')$ $\text{plt.semilogx}()$

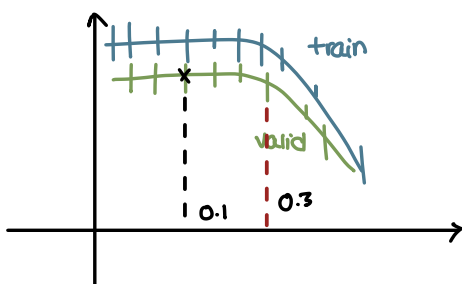
affichage du meilleur point:

$\text{linear_valid_score} = \text{np.mean}(\text{validScore}, \text{axis}=1)$

$\text{bestIndex} = \text{np.argmax}(\text{linear_valid_score})$

$\text{bestNC} = \text{plage-de-valeurs-pour-alpha}[\text{bestIndex}]$

$\text{plt.plot}(\text{bestNC}, \text{linear_valid_score}[\text{bestIndex}], \text{marker}='X', \text{color}='green')$



→ on peut prendre le meilleur validation score

→ si les données sont petites, α contrôle le degré de régularisation,

10^3 : meilleure généralisation (on regularize plus), overfit (espace entre train et valid) est plus petit, barre d'erreur plus petite.

$\text{plt.scatter}(y_test, \text{monModele}.predict(x_test))$

→ x vraie valeur, y valeur prédite

$\text{np.unique}(y_trainval)$

→ entiers

→ classification...

entraînement du modèle sur le meilleur paramètre trouvé:

$\text{monModele} = \text{sklearn.linear_model.Ridge}(\text{alpha}=1000)$

$\text{monModele}.fit(x_trainval, y_trainval)$

$\text{monModele}.score(x_trainval, y_trainval), \text{monModele}.score(x_test, y_test)$

Annale 2022-2023:

```
X, Y = sklearn.datasets.load_digits (n_class = 10, return_X_y = True)
```

```
X_trainval, X_test, Y_trainval, Y_test = sklearn.model_selection.train_test_split(X, Y, test_size = 0.2, random_state = 7)
```

exemple pour $C=1$

```
C = 1
```

```
monModele = sklearn.svm.LinearSVC (C=C, max_iter=1000, dual=False)
```

```
monModele.fit(X_trainval, Y_trainval)
```

```
monModele.score(X_trainval, Y_trainval)
```

faire varier C puis crossfit avec les train et valid score

```
plage_de_valeurs_pour_C = np.logspace(-6, 2, num=9) # entre  $10^{-6}$  et  $10^2$ 
```

```
nbplis CV = 5
```

```
trainScore = np.zeros((plage_de_valeurs_pour_C.shape[0], nbplis(V)))
```

```
validScore = _____
```

```
for i, C in enumerate (plage_de_valeurs_pour_C):
```

```
    mon_modele = sklearn.svm.LinearSVC(C=C, max_iter=1000, dual=False)
```

```
    myCV = sklearn.model_selection.cross_validate(monModele, X_trainval, Y_trainval, cv=nbplisCV, return_train_score=True)
```

```
    trainScore[i] = myCV['train_score']
```

```
    validScore[i] = myCV['test_score']
```

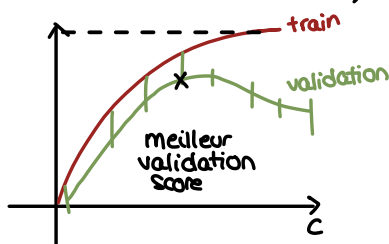
plot

```
plt.figure()
```

```
plt.errorbar(plage_de_valeurs_pour_C, np.mean(trainScore, axis=1), label="train score", yerr=np.std(trainScore, axis=1))
```

```
plt.errorbar(plage_de_valeurs_pour_C, np.mean(validScore, axis=1), label="valid score", yerr=np.std(validScore, axis=1))
```

```
plt.axhline(1, c='k', ls='--') + # affichage du meilleur point
```



intervalle attendue:

$C_{choisi} = 10^{-3}$

$validScore[plage_de_valeurs_pour_C == C_{choisi}]$

$(validScore[plage_de_valeurs_pour_C == C_{choisi}].min(),$

$validScore[plage_de_valeurs_pour_C == C_{choisi}].max())$

```
C = C_choisi
```

```
monModele = sklearn.svm.LinearSVC(C=C, max_iter=1000, dual=False)
```

```
monModele.fit(X_trainval, Y_trainval)
```

```
monModele.score(X_trainval, Y_trainval), monModele.score(X_test, Y_test)
```

* gradient cheat sheet:

let $X \in \mathbb{R}^{N \times D}$, $Y \in \mathbb{R}^N$, $w \in \mathbb{R}^D$, $f_w(X) = Xw$, $N = X.shape[0]$

1) MSE: $J = \frac{1}{2N} \|Xw - Y\|^2$

$\nabla_w J = \frac{1}{N} X^T (Xw - Y)$

matrix: grad = $(X.T @ (X @ w - Y)) / N$ (N,1)

mean: grad = $np.mean((X @ w - Y).reshape(-1, 1) * X, axis=0)$

2) Ridge: $J = \frac{1}{2N} \|Xw - Y\|^2 + \frac{\alpha}{2} \|w\|^2$

$\nabla_w J = \frac{1}{N} X^T (Xw - Y) + \alpha w$

matrix: grad = $(X.T @ (X @ w - Y)) / N + \alpha * w$

mean: grad = $np.mean((X @ w - Y).reshape(-1, 1) * X, axis=0) + \alpha * w$

3) MAE: $J = \frac{1}{N} \sum_i |Xw - Y|_i$

$\nabla_w J = \frac{1}{N} X^T \text{sign}(Xw - Y)$

Matrix: grad = $(X.T @ np.sign(X @ w - Y)) / N$

mean: grad = $np.mean(np.sign(X @ w - Y).reshape(-1, 1) * X, axis=0)$

4) Logistic (BCE): $J = -\frac{1}{N} \sum_i [y_i \log \sigma_i + (1 - y_i) \log (1 - \sigma_i)]$

$\sigma(Xw) = \frac{1}{1 + e^{-(Xw)}}$

$\nabla_w J = \frac{1}{N} X^T (\sigma(Xw) - Y)$

matrix: $y_pred = 1 / (1 + np.exp(-(X @ w)))$

grad = $(X.T @ (y_pred - Y)) / N$

mean: grad = $np.mean((y_pred - Y).reshape(-1, 1) * X, axis=0)$

$$5) \text{ Lasso}(L_1): J = \frac{1}{2N} \|Xw - Y\|^2 + \alpha \|w\|_1$$

$$\nabla_w J = \frac{1}{N} X^T (Xw - Y) + \alpha \text{sign}(w)$$

$$\text{Matrix: grad} = (X.T @ (X @ w - Y)) / N + \alpha * \text{np.sign}(w)$$

6) Huber(δ)

$$r = Xw - Y \quad g_i = \begin{cases} r_i & |r_i| < \delta \\ \delta \text{sign}(r_i) & \text{sinon} \end{cases} \quad \nabla_w J = \frac{1}{N} X^T g$$

$$\text{Matrix: grad} = (X.T @ g) / N \quad \text{Mean: grad} = \text{np.mean}(g.\text{reshape}(-1,1) * X, \text{axis}=0)$$

Partie IV : exemples d'erreurs dans le code.

Annale 2024-2025:

clf = sklearn.svm.SVC(C=1, Kernel='poly', degree=4)

clf.fit(X_train_transformed, Y_train)

train_score = clf.score(X_train_transformed, Y_train)

clf.fit(X_valid_transformed, Y_valid) X il faut pas faire fit deux fois

+ il faut pas fit avec les données de validation

Annale 2022-2023:

preProc = sklearn.decomposition.PCA(n_components=10)

preProc.fit(X_train)

X_train_transformed = preProc.transform(X_train)

X_valid_transformed = preProc.transform(X_valid)

clf = sklearn.svm.SVC(C=1, Kernel='poly', degree=3)

clf.fit(X_train_transformed, Y_train) il faut entraîner le modèle sur les données transformées

$\Rightarrow$ clf.fit(X_train_transformed, Y_train)

$\leadsto$ l'erreur était que on entraînait sur toutes les données (train+val) donc il n'y avait pas en fait de validation set.

Donc la performance de "validation" était la même qu'en train. Avec un split, comme attendu, la perf en validation est légèrement moins bonne.

NOTES: bien transformer le valid comme le train, only fit train, ne pas fit sur les données valides.

autres exemples:

type d'erreur	code erroné	correction	conséquence / explication
1. fit du modèle sur tout le jeu	clf.fit(X, Y)	clf.fit(X_train, Y_train)	Data Leakage: le modèle "voit" les données de test $\leadsto$ score irréaliste (trop bon)
2. fit du transformeur sur tout X	pca.fit(X)	pca.fit(X_train) puis pca.transform(X_valid) / X_test	fuite d'information: les composantes principales sont influencées par le test $\leadsto$ biais dans la validation
3. Mauvais ordre fit/transform	X_train = pca.transform(X_train) pca.fit(X_train)	pca.fit(X_train) X_train = pca.transform(X_train)	erreur logique: le transform précède le calcul des composantes $\leadsto$ résultats incohérents
4. Double fit du modèle	clf.fit(X_train_val) puis clf.fit(X_train_val)	faire un seul fit	le second fit écrase le 1 ^{er} $\leadsto$ incohérence, ou apprentissage sur le mauvais split
5. Evaluation biaisée	clf.score(X_train_val, Y_train_val)	clf.score(X_test, Y_test)	test non indépendant $\leadsto$ on mesure la performance sur des données déjà vues
6. Transform après fit global	pca.fit_transform(X) puis split train/test	split avant, puis fit-transform sur train	Les infos du test influencent la base du modèle $\leadsto$ fuite de données
7. Transform seulement sur le train	pca.fit(X_train) mais X_test inchangé	X_test = pca.transform(X_test)	Le modèle apprend dans un espace transformé. Δ ! les données de validation / test sont dans un espace différent $\leadsto$ prédictions incohérentes

* train dans cross-validate