

Examen de contrôle continu

Durée 1h - Aucun document autorisé
Vendredi 18 octobre 2024

Exercice 1 (6 pts)

Décrire et commenter les résultats obtenus à l'aide du code Python ci-dessous.

Aide paramètre C de la fonction LogisticRegression de sklearn

C : float, default=1.0

Inverse of regularization strength; must be a positive float.

Like in support vector machines, smaller values specify stronger regularization.

```
from sklearn.linear_model import LogisticRegression
from sklearn.preprocessing import StandardScaler
from sklearn.datasets import make_classification

X, y = make_classification(n_samples=300, n_features=40,
                          n_classes=2, flip_y=0.1, class_sep=0.5, random_state=42)

X = StandardScaler().fit_transform(X)

model1 = LogisticRegression(penalty='l1', C=1, class_weight=None,
                             dual=False, solver='saga', fit_intercept=True, l1_ratio=None, max_iter=10000, tol=0.0001)
model1.fit(X, y)
print(model1.score(X, y))
0.8233333333333334

model2 = LogisticRegression(penalty='l1', C=0.1, class_weight=None,
                             dual=False, solver='saga', fit_intercept=True, l1_ratio=None, max_iter=10000, tol=0.0001)
model2.fit(X, y)
print(model2.score(X, y))
0.7866666666666666

from sklearn.model_selection import train_test_split
X, X_test, y, y_test = train_test_split(X, y, test_size=0.25)

model1 = LogisticRegression(penalty=None, class_weight=None,
                             dual=False, solver='saga', fit_intercept=True, l1_ratio=None, max_iter=10000, tol=0.0001)
model1.fit(X, y)
print(model1.score(X_test, y_test))
0.68

model2 = LogisticRegression(penalty='l1', C=0.1, class_weight=None,
                             dual=False, solver='saga', fit_intercept=True, l1_ratio=None, max_iter=10000, tol=0.0001)
model2.fit(X, y)
print(model2.score(X_test, y_test))
0.7466666666666667
```

Exercice 2 (5 pts)

Nous considérons n variables aléatoires indépendantes $y_1, \dots, y_n$ telles que y_i est distribuée suivant une loi binomiale de paramètre $(100, p_i)$. Pour tout $i = 1, \dots, n$, nous supposons que $\log(-\log(1-p_i)) = x_i^T \beta$ où $x_i \in \mathbb{R}^p$ est un vecteur supposé connu et β un vecteur de paramètres inconnus. Montrer qu'il s'agit d'un modèle linéaire généralisé. La fonction de lien canonique a-t-elle été utilisée ?

Exercice 3 (4 pts)

Décrire et commenter les résultats obtenus à l'aide du code Python ci-dessous.

```
from sklearn.model_selection import cross_val_score
from sklearn.linear_model import LinearRegression
from sklearn.linear_model import Lasso
from sklearn.preprocessing import StandardScaler
from sklearn.datasets import make_regression
import numpy as np

X, y = make_regression(n_samples=100, n_features=80, noise=0.1, random_state=42)

scaler = StandardScaler()
X_normalized = scaler.fit_transform(X)

model1 = LinearRegression()
scores1 = cross_val_score(model1, X_normalized, y, cv=5, scoring='neg_mean_absolute_error')
print(scores1)
print(np.mean(scores1))

[-10.96963938 -44.27345903 -40.90757601 -32.031635    -18.96821433]
-29.430104749593802

alpha = 0.1
model2 = Lasso(alpha=alpha)
scores2 = cross_val_score(model2, X_normalized, y, cv=5, scoring='neg_mean_absolute_error')
print(scores2)
print(np.mean(scores2))

[-0.34950166 -0.3056262  -0.30595072 -0.46432015 -0.29677541]
-0.3444348268695149
```

Exercice 4 (5 pts)

1 (2,5 pts) Quelle est la différence entre les modèles logit et probit ?

2 (2,5 pts) Quelle est l'objectif des méthodes de régularisation en regression ?

Correction Kontrollklausur 3

HAX 912X

25/10/2024

(1)

Exercise 1

$$\Rightarrow f(y, \pi) = \binom{h+y-1}{y} \pi^h (1-\pi)^y \prod_{i=1}^y \pi(y)$$

$$f(y; \pi) = \exp\{h \log(\pi) + y \log(1-\pi)\} \prod_{i=1}^y \pi(y)$$

$$\theta = \log(1-\pi) \Rightarrow \pi = 1 - e^\theta$$

$$t(\theta) = -h \log(1 - e^\theta)$$

$$t'(\theta) = \frac{h e^\theta}{1 - e^\theta} \quad \text{Formlich exponentiell 1-wahrsch}$$

$$\log(1-\pi) = \beta_1 + \beta_2 \kappa$$

$$\pi E(y) = t'(\theta) = \frac{h e^\theta}{1 - e^\theta} = h \left(\frac{1-\pi}{\pi} \right)$$

$$\pi = 1 - e^{\beta_1 + \beta_2 x}$$

(2)

$$\Rightarrow E(y) = h \left\{ \frac{e^{\beta_1 + \beta_2 x}}{1 - e^{\beta_1 + \beta_2 x}} \right\}$$

C'est bien un modèle linéaire généralisé.

En fait, $0 = \log(1 - \pi) = \beta_1 + \beta_2 x$
donc c'est bien le lien canonique qui a été utilisé.

$$2) \pi = 1 - e^{\beta_1 + \beta_2 x}$$

$$\angle V_M(\beta_1, \beta_2) = h \sum_{i=1}^M \log(1 - e^{\beta_1 + \beta_2 x_i})$$

$$+ \sum_{i=1}^M y_i (\beta_1 + \beta_2 x_i) + C$$

$$\frac{\partial V_M(\beta_1, \beta_2)}{\partial \beta_1} = -h \sum_{i=1}^M \left[\frac{e^{\beta_1 + \beta_2 x_i}}{1 - e^{\beta_1 + \beta_2 x_i}} \right] + \sum_{i=1}^M y_i$$

$$\frac{d\Delta V_m}{d\beta_2}(\beta_1, \beta_2) = -k \sum_{i=1}^m \alpha_i \left[\frac{e^{\beta_1 + \beta_2 \alpha_i}}{1 - e^{\beta_1 + \beta_2 \alpha_i}} \right] + \sum_{i=1}^m y_i \alpha_i \quad (3)$$

Exercice 2

Problème de régression logistique
 $\rightarrow y \in \{0, 1\}$ = variable indépendante
 la présence ou l'absence d'un
 défaut, variable à expliquer
 $\rightarrow x \in \mathbb{R}$ Température au moment
 du bon ciment, variable explicative

$$P(y=1|x) = \frac{e^{\beta_0 + \beta_1 x}}{1 + e^{\beta_0 + \beta_1 x}}$$

On estime (β_0, β_1) avec les 24
 données disponibles, on obtient
 $(\hat{\beta}_0, \hat{\beta}_1)$.

④

Enfin, on veut

$P(y=1 | x=31)$ avec

$$\left[\frac{e^{\hat{\beta}_0 + 31 \hat{\beta}_1}}{1 + e^{\hat{\beta}_0 + 31 \hat{\beta}_1}} \right]$$

avec avec ça =

model <- glm (age ~ temp, family =
"binomial", data = challenges)

coeff <- model\$coefficients

a <- exp (coeff [1] + coeff [2] * 31)

a / (1 + a)

Exercice 3

(5)
1 point $\rightarrow$ les densités
montrent que la variable
prix est intéressante pour
expliquer le revenu et
l'investissement, mais que ce n'est
pas le cas de la variable
Storm Duration.

1 point $\rightarrow$ un premier modèle
logistique est mis en œuvre
Prix est variable tout asymptotique
et modèle, le coefficient a
un intervalle très petit.

1 point $\rightarrow$ un deuxième modèle
logistique est mis en œuvre
Storm Duration est une variable
et l'intervalle très (très) petit.

1 point $\rightarrow$ pour le module 1
avec le clip, on estime par
validation overfit $\approx$ 10 ensemble
après 50 fois la matrice de
confusion difficile à pratiquer
Chap 2 (Russellment) mais on y
arrive on va.

1 point $\rightarrow$ pour le module 2
avec Storm Distribution, incompa-
rable à pratiquer le Chap 1

Exercice 4

1 point $\rightarrow$ $n = 5000$ individus
 $p = 60$ variables
régression $\rightarrow X_1$
on suppose 100 variables
à l'arrêt $\rightarrow X_2$

2 points $\rightarrow$ 2 merites de $\textcircled{7}$
points relatifs, l'un sur les
variables de bruit, l'autre
avec

champs reportés en classe

- un échantillon de training
- un échantillon de test $_{1000}$ $_{4000}$

2 points $\rightarrow$ comparaison sur
échantillons de test
même résultat
sur une inférence de bruit
