

Scoring Appliqué à la Détection du Risque – Chapitre 2 – Méthodes probabilistes

Contexte et objectif

Les données. Elles sont constituées de n couples $(\mathbf{x}_1, \mathbf{z}_1), \dots, (\mathbf{x}_n, \mathbf{z}_n)$. Pour $i \in \{1, \dots, n\}$: $\mathbf{x}_i = (x_{i,1}, \dots, x_{i,d})' \in \mathcal{X}$ est la valeur des prédicteurs pour l'individu i et $\mathbf{z}_i = (z_{i,1}, \dots, z_{i,G}) \in \{0, 1\}^G$ indique sa classe : $z_{i,g} = 1 \Leftrightarrow \mathbf{x}_i$ appartient à la classe g . L'ensemble $\mathcal{X}$ dépend de la nature des données prédictives (quantitatives continues/discrètes, qualitatives, mixtes etc).

L'objectif. Estimer la classe de $\mathbf{x}_{n+1} \in \mathcal{X}$.

La méthode. Définir G fonctions de score $\mathbf{s}_g : \mathcal{X} \rightarrow \mathbb{R}$ permettant de décider : $\hat{z}_{n+1,g} = 1 \Leftrightarrow \mathbf{s}_g(\mathbf{x}_{n+1}) = \max\{\mathbf{s}_j(\mathbf{x}_{n+1}); j = 1, \dots, G\}$.

Analyse discriminante probabiliste

L'analyse discriminante probabiliste s'inscrit dans une approche générative : elle propose un modèle paramétrique expliquant la genèse des observations $\{(\mathbf{x}_i, \mathbf{z}_i)\}$.

Le modèle. On suppose que :

- les couples $(\mathbf{x}_1, \mathbf{z}_1), \dots, (\mathbf{x}_n, \mathbf{z}_n)$ sont des réalisations indépendantes de $(\mathbf{x}, \mathbf{z})$ où :
 $\mathbf{x} = (x_1, \dots, x_d)' \in \mathcal{X}$ est un vecteur aléatoire de prédicteurs,
 $\mathbf{z} = (z_1, \dots, z_G)' \in \{0, 1\}^G$ est un vecteur aléatoire donnant la classe de $\mathbf{x}$ ($\sum_{g=1}^G z_g = 1$),
- $\mathbf{z}$ est distribué selon une loi multinomiale d'ordre G à un essai de paramètre : $(\pi_1, \dots, \pi_G)$; $0 < \pi_g$ et $\sum_{g=1}^G \pi_g = 1$,
- le vecteur conditionnel $(\mathbf{x} | z_g = 1)$ est distribué selon une loi à support dans $\mathcal{X}$, de fonction de probabilité : $f_g(\bullet; \theta_g)$.

Les fonctions $f_g(\bullet; \theta_g)$ appartiennent généralement à une même famille de lois qui dépend de la nature des données.

La méthode.

- estimation du paramètre $\theta = (\pi_1, \dots, \pi_G, \theta_1, \dots, \theta_G)$ par maximum de vraisemblance :

$$\hat{\theta} = \operatorname{argmax}_{\theta} \left\{ \ell(\theta; \{\mathbf{x}_i, \mathbf{z}_i\}) = \log \left\{ \prod_{i=1}^n \prod_{g=1}^G [\pi_g f_g(\mathbf{x}_i; \theta_g)]^{z_{i,g}} \right\} \right\}$$

- estimation des probabilités conditionnelles d'appartenance aux classes :

$$t_g(\mathbf{x}_{n+1}) = E(z_{n+1,g} | \mathbf{x}_{n+1}) = \hat{\pi}_g f_g(\mathbf{x}_{n+1}; \hat{\theta}_g) / \sum_{j=1}^G \hat{\pi}_j f_j(\mathbf{x}_{n+1}; \hat{\theta}_j)$$

- prédiction : en supposant les coûts de mauvais classement homogènes, la règle de Bayes visant à minimiser le coût moyen d'une erreur de classement revient à affecter $\mathbf{x}_{n+1}$ à la classe correspondant à la plus grande valeur des probabilités conditionnelles : $t_g(\mathbf{x}_{n+1})$; $g = 1, \dots, G$ (règle du MAP^a).

Choix d'un modèle/d'un classifieur.

Les fonctions de probabilité associées aux composantes : $f_g(\bullet, \theta_g)$ déterminent entièrement le modèle/classifieur. Plusieurs modèles peuvent être en concurrence ; on les départage alors en s'appuyant sur (au choix) :

- (i) un critère de vraisemblance comme : $AIC = \ell(\hat{\theta}; \{\mathbf{x}_i, \mathbf{z}_i\}) - \eta$, $BIC = \ell(\hat{\theta}; \{\mathbf{x}_i, \mathbf{z}_i\}) - (\eta/2) \times \log n$, où η est la dimension du paramètre
- (ii) l'erreur de classement : $\epsilon = \sum_{g=1}^G \hat{\pi}_g \int_{\{\mathbf{x} \in \mathcal{X}; \hat{z}_g \neq 1\}} f_g(\mathbf{x}; \hat{\theta}_g) d\mathbf{x}$ estimée par : resubstitution, validation croisée, ou bootstrap.

Les critères de vraisemblance (i) favorisent l'estimation de densité tandis que (ii) donne une importance prépondérante à la classification.

Remarque. En analyse discriminante, les fonctions de score sont les probabilités conditionnelles : $\mathbf{s}_g(\mathbf{x}) = E(z_g | \mathbf{x})$.

L'analyse discriminante en pratique.

L'affectation de $\mathbf{x}_{n+1}$ à une classe résulte du modèle inféré. Or le modèle dépend du contexte : on prendra pour fonction $f_g(\bullet; \theta_g)$ la densité d'une loi normale par exemple lorsque les observations sont continues, la fonction de probabilité d'une loi multinomiale lorsque les observations sont qualitatives, etc. Il existe des logiciels permettant d'inférer un modèle dans la plupart des cas classiques : données continues, discrètes, nominales, mixtes, etc. Table 1 cite quelques uns de ces logiciels.

A. Lourme, Faculté d'économie, gestion & AES, Université de Bordeaux <http://alexandre.lourme.free.fr>

^aMaximum A Posteriori

type de données	contexte	lois conditionnelles ^b	package	références	fonctions utiles
quantitatives continues	données standard	normales	Rmixmod	[7, 8]	mixmodLearn, mixmodPredict
	données bruitées	Student	teigen	[2, 1]	teigen
	données de grande dimension	normales	HDclassif	[3, 4]	hdda
discrètes	données de comptage	Poisson	pgmm	[9]	pgmmEM
qualitatives	données nominales	multinomiales	MixtComp	[5]	Learn, Predict
mixtes	données quantitatives, qualitatives, ...	mixtes	Rmixmod	[7, 8]	mixmodLearn, mixmodPredict
			MixtComp	[5]	Learn, Predict

Table 1: Les logiciels dédiés à l'inférence d'un modèle d'analyse discriminante selon le contexte (données continues/qualitatives/mixtes/etc.)

Régression logistique

La régression logistique s'inscrit dans une approche prédictive : elle propose un modèle semi paramétrique sur $\mathbf{z}$ conditionnellement à $\mathbf{x}$.

Le modèle. On se place dans le cas de données continues ($\mathcal{X} = \mathbb{R}^d$) et on suppose la partition binaire^c ($G = 2$). La régression logistique modélise seulement la probabilité conditionnelle $t_1(\mathbf{x})$ d'appartenir à la classe 1 : $t_1(\mathbf{x}) = \exp(\alpha + \beta' \mathbf{x}) / (1 + \exp(\alpha + \beta' \mathbf{x}))$ où $\alpha \in \mathbb{R}$ et $\beta \in \mathbb{R}^d$. Ainsi $\log\{t_1(\mathbf{x}) / (1 - t_1(\mathbf{x}))\} = \alpha + \beta' \mathbf{x}$ et la frontière séparant les deux classes est un hyperplan.

La méthode.

◦ estimation du paramètre $\theta = (\alpha, \beta)$ par maximum de vraisemblance conditionnelle :

$$\hat{\theta} = \operatorname{argmax}_{\theta} \left\{ \sum_{i=1}^n z_{i,1}(\alpha + \beta' \mathbf{x}_i) - \log(1 + \exp(\alpha + \beta' \mathbf{x}_i)) \right\}. \quad (1)$$

Remarque. Les équations de vraisemblance permettant de déterminer (1) ne sont pas linéaires ; un algorithme de Newton-Raphson permet cependant de déterminer $\hat{\theta}$. En pratique, sous R, l'estimation de θ peut être réalisée grâce à la fonction `glm`.

◦ prédiction :

- soit la fonction de score est : $\mathbf{s}_1(\mathbf{x}) = \hat{\alpha} + \hat{\beta}' \mathbf{x}$ et $\mathbf{x}$ est affecté à la classe 1 si $\mathbf{s}_1(\mathbf{x}) > 0$, à la classe 2 dans le cas contraire
- soit la fonction de score est : $\mathbf{s}_1(\mathbf{x}) = t_1(\mathbf{x}; \hat{\alpha}, \hat{\beta}) = \exp(\hat{\alpha} + \hat{\beta}' \mathbf{x}) / (1 + \exp(\hat{\alpha} + \hat{\beta}' \mathbf{x}))$ et $\mathbf{x}$ est affecté à la classe 1 si $\mathbf{s}_1(\mathbf{x}) > 0,5$, à la classe 2 dans le cas contraire.

References

- [1] Jeffrey L Andrews and Paul D McNicholas. Model-based clustering, classification, and discriminant analysis via mixtures of multivariate t-distributions. *Statistics and Computing*, 22(5):1021–1029, 2012.
- [2] Jeffrey L Andrews, Paul D McNicholas, and Maintainer Jeffrey L Andrews. Package ‘teigen’. 2015.
- [3] Laurent Bergé, Charles Bouveyron, and Stéphane Girard. HDclassif: An R package for model-based clustering and discriminant analysis of high-dimensional data. *Journal of Statistical Software*, 46(6):1–29, 2012.
- [4] Charles BOUYEYRON. Hdclassif: an r package for model-based classification of high-dimensional data. *Talk, Université Paris*, 1, 2012.
- [5] MODAL Team INRIA Lille Nord Europe. MixtComp MODELS for Data Analysis and Learning. <https://modal.lille.inria.fr/wikimodal/doku.php?id=mixtcomp>, 2015.
- [6] Gérard Govaert. *Data analysis*. John Wiley & Sons, 2013.
- [7] Rémi Lebre, Serge Iovleff, Florent Langrognet, Christophe Biernacki, Gilles Celeux, and Gérard Govaert. Rmixmod: the r package of the model-based unsupervised, supervised and semi-supervised classification mixmod library. *Journal of Statistical Software*, pages In–press, 2015.
- [8] Remi Lebre, Serge Iovleff, Florent Langrognet, Maintainer Remi Lebre, and LinkingTo Rcpp. Package ‘rmixmod’. 2012.
- [9] PD McNicholas, KR Jampani, AF McDaid, TB Murphy, and L Banks. pgmm: Parsimonious gaussian mixture models. *R package version*, 1:24–47, 2011.

^cil est possible de généraliser la régression logistique à plus de deux classes (voir [6], p.194)