



Efficacité des modèles de classification

par Louis Rossouw, Gen Re, Le Cap

Les assureurs développent de plus en plus de modèles de prédiction qu'ils utilisent dans leurs processus d'assurance. Ces modèles reposent souvent sur des techniques traditionnelles, mais les techniques d'apprentissage automatique sont de plus en plus appliquées.

Dans la pratique, ces techniques sont appliquées aux cas de sélection des risques, aux demandes de souscription de polices, voire aux déclarations de sinistre. Un exemple peut prendre la forme d'un modèle qui est utilisé pour prédire les cas qui seront évalués comme étant « standard » avant qu'un souscripteur ne s'en aperçoive. C'est ce que l'on appelle un « modèle de classification » utilisé pour classer les données dans des buckets discrets (oui ou non, standard ou non, etc.).

Les techniques de modélisation utilisées dans ces applications peuvent être relativement simples ou bien très complexes. Voici quelques exemples de techniques :

- Régression logique (généralement un modèle linéaire généralisé – MLG)
- Arbres de décision
- Forêts aléatoires
- Machines à vecteurs de support
- Techniques de descente du gradient
- Réseaux de neurones artificiels

Les techniques traditionnelles, telles que des modèles basés sur la régression, produisent des modèles qui sont lisibles par des humains. L'impact de chaque variable que comprend le modèle est clairement visible sur le résultat. La plupart des techniques d'apprentissage automatique produisent des modèles que l'observation directe ne permet pas de comprendre facilement. Ils produisent des résultats, mais les mécanismes internes sont cachés ou trop complexes pour pouvoir parfaitement les comprendre. Ce sont les modèles dits de « boîte noire ».

Avec un tel choix de modèles, comment évaluer la précision et la qualité de ces modèles pour identifier le meilleur ? Comment tenons-nous compte de l'efficacité

Contenu

Données d'entraînement et de tests	2
La matrice de confusion	2
Scores des modèles et seuil	3
Fonction d'efficacité du récepteur ou courbe ROC	3
Coefficient de Gini	4
Comparaison des modèles	4
Modèles d'ensemble	4
Optimisation des affaires	4
Conclusion	5

des modèles faciles à comprendre, par rapport aux différents modèles de boîte noire ? Comment évaluons-nous la valeur relative pour l'assureur de ces modèles ?

Dans la suite de l'article, nous nous concentrerons sur un classifieur binaire simple capable de prédire si une demande de souscription doit être considérée comme standard (sans supplément ni condition de sélection des risques) ou non (refusée, avec une surprime ou clause d'exclusion). Ce modèle peut être utilisé pour passer outre la sélection des risques traditionnelle afin de gagner du temps et faire des économies pour un sous-ensemble de cas.

Données d'entraînement et de tests

Pour créer un modèle classifieur, on a généralement un ensemble de données avec des cas historiques et le résultat enregistré de ces cas. Dans notre exemple de sélection des risques, il peut s'agir de données concernant le demandeur et la décision enregistrée (standard ou non).

Il est préférable de répartir les données passées dans au moins deux catégories :

- Les données d'entraînement qui seront utilisées pour intégrer le(s) modèle(s) en question ; ces données sont utilisées pour « entraîner » le modèle.
- Les données de test qui seront utilisées pour évaluer le(s) modèle(s) afin de déterminer l'efficacité du/des modèle(s)

Les données de validation sont également utilisées très souvent pour affiner les paramètres de modélisation avant de tester le modèle.

La raison principale de la séparation entre les tests et l'entraînement est de s'assurer que le modèle fonctionne bien avec les données sur lesquelles il n'a pas été entraîné. Cela permet notamment d'identifier le problème de surapprentissage ou surajustement, qui se produit lorsqu'un modèle donné semble donner des prédictions trop précises à partir des données d'entraînement. Lorsque ce modèle est ensuite vérifié par rapport aux autres données de test, sa performance se dégrade toutefois considérablement. On constate alors que le modèle « surapprend » les données d'entraînement.

Entre 10 et 30 % des données sont retenues comme données de test. Cela dépend de la disponibilité globale des données et du degré de surapprentissage possible des modèles. Il est possible de sélectionner des données de test comme sous-ensemble aléatoire de vos données ou, par exemple, baser la sélection des données sur le temps (en sélectionnant la dernière année de données comme sous-ensemble de test).

La matrice de confusion

Pour évaluer la qualité d'un classifieur binaire, il est possible de générer une matrice de confusion avec nos données de test. Il suffirait d'appliquer le modèle sur ces données afin de produire des résultats prédits, et intégrer aussi les résultats réels dans les données. Puis de compiler une matrice avec les nombres de cas pour lesquels le modèle a identifié le résultat comme positif (standard dans notre exemple) et le résultat réel était positif (également standard). C'est le nombre de « vrais positifs ». Nous comptons également le nombre de cas pour lesquels le modèle a correctement prédit le négatif comme des « Vrais négatifs ». Nous comptons également les erreurs lorsque le modèle a prédit un négatif mais le résultat était positif et inversement. La matrice de confusion ci-dessous est ainsi générée :

		Prédit	
		Positif	Négatif
Réal	Positif	Vrais positifs	Faux négatifs
	Négatif	Faux positifs	Vrais négatifs

Dans une matrice de confusion, les faux positifs sont également des erreurs de type I et les faux négatifs sont considérés comme des erreurs de type II.

On peut ensuite utiliser cette matrice de confusion pour classer le modèle à l'aide de divers paramètres :

$$\begin{aligned} \text{Précision} &= \frac{\text{Faux positifs} + \text{vrais négatifs}}{\text{Nombre total de cas}} \\ \text{Sensibilité} &= \frac{\text{Vrais positifs}}{\text{Vrai Positifs} + \text{Faux négatifs}} = \frac{\text{Vrais Positifs}}{\text{Réellement positifs}} \\ \text{Spécificité} &= \frac{\text{Vrais négatifs}}{\text{Vrais négatifs} + \text{Faux positifs}} = \frac{\text{Vrais négatifs}}{\text{Réellement négatifs}} \end{aligned}$$

La mesure de précision est une mesure globale de la précision. La sensibilité indique le nombre de positifs qui sont réellement identifiés comme tels. L'interaction de ces variables indique l'efficacité

d'un modèle. Par exemple, si nous avons deux modèles prédisant si des cas sont standard ou non, nous pouvons tabuler une matrice de confusion pour chaque modèle.

Le modèle A est un modèle très simple (et très imprécis) qui prédit que chaque cas sera standard.

Modèle A		Prédit	
		Standard	Non-standard
Réal	Standard	80	0
	Non-standard	20	0

D'après la matrice de confusion, nous pouvons estimer qu'il y a eu 100 cas. Notre modèle prédit que tous les cas seront standard. Notre sensibilité est alors de 100 % et notre précision semble satisfaisante à 80 %. Cependant, avec une spécificité à 0 %, cela indique que le modèle est peu performant.

Il convient de prendre en compte plusieurs mesures pour connaître l'efficacité d'un modèle. Un modèle plus réaliste ressemblerait à celui reproduit ci-dessous (sur les mêmes données) :

Modèle B		Prédit	
		Standard	Non-standard
Réal	Standard	61	19
	Non-standard	8	12

Dans ce cas, la précision globale est de 73 %, la sensibilité de 76,25 % et la spécificité de 60 %. Ce modèle semble raisonnable.

Scores des modèles et seuil

La plupart des modèles classifieurs ne produisent pas qu'une classification binaire comme résultat. Ils calculent généralement un score pour classer les cas comme positifs ou négatifs. Le score est en général converti en pourcentage. Ce score n'implique pas toujours forcément une vraie probabilité, notamment pour les techniques d'apprentissage automatique, et indique souvent qu'un classement des cas, plutôt qu'une probabilité stricte.

Sachant que chaque cas aurait un score, il conviendrait d'attribuer un seuil (ou une limite)

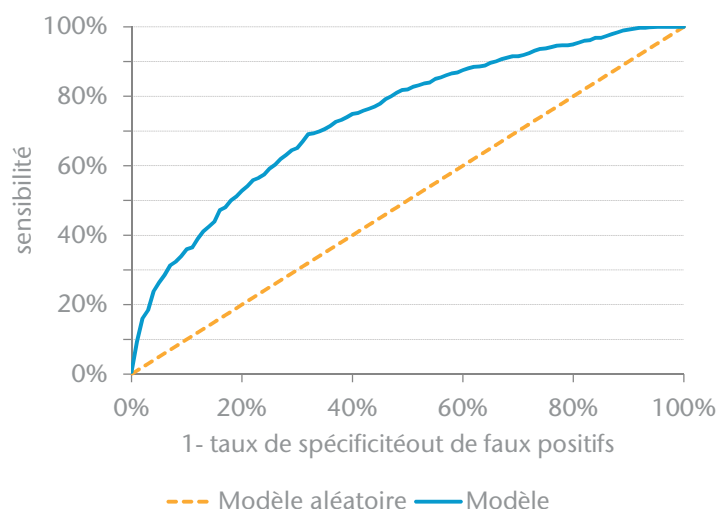
au-delà duquel le résultat du modèle pourrait être considéré comme positif. Par exemple, un modèle pourrait donner différents scores pour différents cas, et avec un seuil de 80 %, traiterait un cas comme un cas standard uniquement lorsque le score est supérieur à 80 %.

Chacun de ces seuils suppose une matrice de confusion spécifique. Avec un seuil de 0 %, nous finirions avec un modèle A représenté ci-dessus, nous prédirions chaque cas comme standard, avec la sensibilité de 100 % mais une spécificité de 0 %. De la même manière, avec un seuil de 100 %, tous les cas seront classés comme non-standard, avec une sensibilité de 0 % et une spécificité de 100 %.

Fonction d'efficacité du récepteur ou courbe ROC

La modification du seuil d'un modèle donné sur toutes les valeurs entre 0 et 100 % permet de tracer une courbe des différentes valeurs de spécificité et de sensibilité. La courbe ROC trace la sensibilité (axe Y) par rapport à la spécificité (axe X) pour chaque valeur du seuil. Notez que la spécificité (1) est également le taux de faux positifs. La figure 1 est un exemple de cette courbe ROC.

Figure 1 – Courbe caractéristique du récepteur



En bas à gauche, on retrouve le cas où le seuil est de 100 %. C'est l'équivalent de notre modèle pour prédire tous les cas comme étant non-standard (sensibilité de 0 % mais spécificité de 100 % ou taux de faux positifs est de 0 %). En haut à droite, nous avons le cas de figure dans lequel tous les cas sont prédits comme standard (sensibilité de 100 % mais spécificité de 0 % ou taux de faux positifs est de 100 %).

Un nouvel échantillon de données de validation (non utilisé pour l'entraînement ou des tests) serait généralement retenu pour produire ce type de modèle d'ensemble.

Optimisation des affaires

Pour un modèle avec une courbe ROC donnée, nous pouvons maintenant décider de la manière d'appliquer ce modèle en pratique. Sachant que chaque type d'erreur (faux positifs et faux négatifs) aurait des coûts et des avantages pour l'entreprise, nous pouvons ensuite estimer le point auquel le coût est réduit ou auquel l'avantage est maximisé, fixant ainsi le meilleur seuil pour chaque modèle. Ces résultats peuvent nous permettre de comparer les meilleures performances des différents modèles pour prendre une décision finale sur un modèle en particulier.

Dans l'exemple de la modélisation des décisions de tarification standard, nous pouvons prendre en compte les implications de chaque catégorie de résultat sur la valeur actualisée des profits par demande :

- Les vrais positifs (cas correctement classés par le modèle comme standard) pourraient voir la valeur par police augmenter, car nous constaterions des placements plus élevés de cette catégorie en raison de frais de vente par police plus faibles (taux de conversion plus élevé du fait d'un processus rationalisé pour le client) et des frais de sélection des risques et médicaux plus bas.
- Les faux positifs (cas classés à tort comme standard) peuvent être confrontés à des problèmes d'augmentation des sinistres par rapport aux primes. Il existe également un risque de sélection adverse si les demandeurs savent comment influencer leurs scores.
- Les faux négatifs (cas qui sont classés à tort comme non-standard) peuvent être très proches du processus actuel. Il se peut donc que nous devions utiliser un taux de conversion actuel et des taux d'utilisation pour évaluer la valeur actualisée des profits par demande (en supposant que l'underwriting suive notre approche actuelle).
- Les vrais négatifs (cas correctement classés comme négatifs) peuvent encore une fois être cohérents avec notre traitement actuel des cas non-standard.

Au vu des valeurs pour chaque élément ci-dessus, il est possible d'estimer un seuil optimal pour notre modèle, seuil qui apportera une valeur maximale par demande à l'entreprise. Cela correspondrait à un point unique sur la courbe ROC. La sensibilité du résultat sélectionné doit être testée par rapport aux hypothèses car la plupart de ces dernières, établies pour déterminer la valeur des différents scénarios, pourraient être subjectives.

Si nous disposons de plusieurs modèles (où les courbes ROC se croisent comme précédemment), ou un modèle pour lequel nous ne pouvons pas calculer la courbe ROC, nous pouvons alors comparer la meilleure valeur produite à partir des différents modèles pour déterminer le meilleur d'entre eux. Pour comparer ces valeurs, il convient de s'assurer que le coût de fonctionnement du modèle est inclus. Certains modèles peuvent avoir des implications en matière de coûts, soit du fait des données utilisées, soit en raison de problèmes d'application techniques.

Conclusion

Cet article a présenté certaines approches permettant d'évaluer la performance des modèles utilisés dans le cadre de problèmes de classification.

Nous avons apporté les éclaircissements suivants :

- Il est nécessaire d'avoir des échantillons de données pour les tests.
- Il est possible de comparer la performance de ces modèles en utilisant des mesures dans une matrice de confusion.
- Il est possible d'estimer certaines données, comme la spécificité et la sensibilité à partir de ces matrices de confusion.
- On peut également envisager d'utiliser la courbe ROC et la zone sous la courbe pour se faire une idée de la qualité du modèle.
- Des techniques d'ensemble peuvent être utilisées pour combiner des scores de différents modèles pour produire de meilleurs modèles.
- Une fois que nous appliquons des valeurs commerciales attendues pour différents résultats du modèle, il peut être plus simple de déterminer la valeur qui représenterait le meilleur compromis entre la sensibilité et la spécificité pour l'entreprise.

- Il peut y avoir des considérations pratiques à prendre en compte, notamment des hypothèses subjectives dans l'évaluation de divers résultats et des détails techniques de mise en œuvre qui peuvent donner lieu à des résultats différents.
- La modification du processus peut entraîner des changements de comportement, ce qui invalide la modélisation. Dans ce contexte, il convient d'évaluer les comportements particulièrement anti-sélectifs.

Cette vue d'ensemble relativement simple donne une idée de la manière d'évaluer ces modèles objectivement et de la façon d'évaluer la valeur du modèle pour l'entreprise en tenant compte de sa performance.

Références :

James, G., Witten, D., Hasties, T. & Tibshirani, R. (2013). An Introduction to Statistical Learning. Springer.

Matloff, N. (2017). Statistical Regression and Classification: From Linear Models to Machine Learning. CRC Press.

Fawcett, T. (2006) An introduction to ROC analysis. Pattern Recognition Letters 27, 861–874.

L'auteur

Louis Rossouw occupe les postes d'actuaire analyste et de chargé d'études au sein du bureau de Gen Re au Cap pour lequel il est chargé des marchés britannique et sud-africain. Louis a rejoint la société en 2001, après avoir travaillé sur la mise au point des barèmes individuels et sur le développement de produits, la tarification de groupe et les provisions. Il a également passé deux ans au sein de la filiale de Gen Re à Singapour en tant qu'actuaire régional en chef. Il peut être contacté par tél. au +27 21 412 7712 ou par courriel lrossouw@genre.com.



The difference is...the quality of the promise.

genre.com | genre.com/perspective | Twitter: @Gen_Re

General Reinsurance AG
Theodor-Heuss-Ring 11
50668 Cologne
Tel. +49 221 9738 0
Fax +49 221 9738 494

General Reinsurance AG—Succursale Paris
21, rue Balzac
75008 Paris
Tel. +33 1 5367 7676
Fax +33 1 5367 46464

Editors:
Ulrich Pasdika, ulrich.pasdika@genre.com
Ross Campbell, ross_campbell@genre.com

Photos: © getty images – underworld111, aedkais, Fredex8

General Reinsurance AG 2018

Les articles publiés sont protégés par un copyright. Ceux qui sont signés ne reflètent pas nécessairement l'opinion de la direction de la publication ou de la rédaction. L'ensemble des informations présentées ont fait l'objet de recherches minutieuses et ont été rassemblées au mieux de nos connaissances en la matière. Nous déclinons toutefois toute responsabilité quant à leur exactitude, leur exhaustivité ou leur caractère d'actualité. En particulier, ces informations ne doivent pas être interprétées comme des conseils d'ordre juridique, auxquels elles ne peuvent en aucun cas se substituer.