

Méthode Bootstrap et applications en R

①

Preface :

- en stat, "bootstrapping" renvoie à faire de l'inférence à partir d'une distribution d'échantillonnage d'une statistique en rééchantillonnant l'échantillon lui-même avec remise, comme si nous avions affaire à une population finie.
 - A condition que la distribⁿ rééchantillonnée imite bien l'originale, l'inférence est robuste - la précision s'améliore quand la taille de l'échantillon grandit en vertu du théorème central limite, lorsque celui-ci s'applique.
 - L'idée du rééchantillonnage est apparue dans les travaux de R.A. Fisher (1935) et d'E.J.G. Pitman (1937, 1938).
 - Concernant le bootstrap et le tirage avec remise, en utilisant des approx. Monte Carlo, ce sont les travaux d'Efron (Brad) qui (papier Annals of Stats, 1977).
- Dans le papier d'Efron, le lien du bootstrap avec d'autres techniques de rééchantillonnage est présentée, c'est le début de l'intérêt de la communauté statistique/stats.
(ex: jackknife)
- L'avantage du bootstrap combiné aux approx. Monte Carlo est de pouvoir faire de l'inférence statistique sous hypothèses faibles (ex: hyp. paramétrique, de lois de proba...).

INTRODUCTION

I - Histoire

- Le bootstrap fait ~~partie~~ partie des méthodes de rééchantillonnage (stats non-paramétriques).
- Quelques techniques sont + anciennes que le bootstrap : les méthodes de permutation de Fisher (1935) et Pitman (1937, 1938) ; le jackknife avec Owen (1945).
↳ (John Tukey)

- Le bootstrapping est devenu applicable d'un point de vue pratique depuis l'avènement de l'informatique et le couplage avec les méthodes Monte Carlo.
- A la base, le but du bootstrap était de construire une approximation du jackknife.
- Le Jackknife est utile à la base pour les petits échantillons, car ^{très} consommateur en ressources de calcul.
- Bootstrap ordinaire implique un tirage avec remise n fois d'un échantillon de taille n
 $\Rightarrow n^n$ possible distinct ordered bootstrap samples (bien que certains soient équivalents sous l'hypothèse d'échangeabilité liés aux permutations).
 $\Rightarrow$ l'énumération des $\neq$ échantillons bootstrap est infaisable dès que n devient grand.
- L'idée de prendre des échantillons Monte Carlo de taille n avec remise à partir de l'échantillon d'origine existait déjà (Julian Simon, 1968), et permettait de s'affranchir de l'hypothèse paramétrique d'un modèle de probabilité par la génération de l'échantillon d'origine.
- Efron (1973) a connecté le bootstrapping avec le Jackknife, la méthode Delta, la validation croisée et les tests de permutation. Il a été le 1^{er} à démontrer que le bootstrap était compétitif avec ces autres méthodes pour estimer l'écart-type d'un estimateur.
- Efron a également vite compris que le bootstrap pouvait être utile aussi pour le calcul d'intervalles de confiance, de tests d'hypothèses, ... (voir Efron and Gong (1983), ...)
- Le bootstrap s'est énormément développé dans les années 80 et 90.
- Auteurs connus pour le bootstrap: Efron, Tibshirani, Chernick, Therneau, Michael Martin, Peter Hall, David Freedman.
- Malgré des résultats théoriques prometteurs dans les années 80, quelques exemples montrent que le bootstrap n'est parfois pas convaincant... L'estimateur ne CV pas lors le vrai paramètre pd la taille de l'échantillon devient infini).

- Généralement, le bootstrap est consistant lorsque le TCL s'applique - (2)
- Une condition suffisante est la condition de Lyapunov: existence du $2+\delta^e$ moment de la distribution de la population.
- Il existe des versions du bootstrap pour remédier aux pb d'inconsistance dans certains cas bien connus.

II - Définition et lien avec la méthode Delta et d'autres méthodes de rééchantillonnage

- Définition informelle: le but du bootstrapping est d'estimer un paramètre ^{via Monte Carlo} basé sur des données (ex: moyenne, médiane, écart-type). On s'intéresse aussi aux propriétés de la distribution de l'estimation de ce paramètre; pour construire des intervalles de confiance par exemple. Mais sans faire d'hyp. restrictives sur la ^(forme de la) distribution d'origine des données observées.
- Pour le cas i.i.d., la brique de base est la distribution empirique: c'est la distribution directe qui donne un poids égal à chaque observation (poids = $\frac{1}{n}$) si taille échantillon = n .
Notée F_n en général. (cf schéma)
- La plupart des paramètres que nous étudions (ex: la moyenne d'un échantillon) sont des fonctions de la distribution inconnue F de la population, des fonctions qui prennent F en entrée et renvoient un réel correspondant. Dans notre cas, on s'intéresse aux fonctionnelles des F d.R. Par exemple,
 $\rightarrow$ si $E[X] = \mu$ où $X \sim F$, alors? $\mu = \int x dF(x)$ $\bar{\mu}_n = \int x dF_n(x)$
 $\rightarrow$ si $\text{Var}(X) = \sigma^2$ alors? $\sigma^2 = \int (x - \mu)^2 dF(x)$ $= \frac{1}{n} \sum_{i=1}^n x_i$
- $\Rightarrow$ On note que les estimations données par ces paramètres en pratique (par ex la moyenne) sont issues des mêmes fonctionnelles appliquées à F_n - (la F d.R. empirique donc).
- L'idée du bootstrap est d'utiliser ce que l'on sait à partir des données, sans introduire d'hypothèse non-justifiées à propos de la distribution de (la population - l'échantillon).

- Principe du bootstrap: prend comme F la distribution de la population, et $T(F)$ la fonctionnelle qui définit le paramètre d'intérêt (ex: $T(F) = \mu = \int x dF(x)$).

En estimant ce paramètre depuis un échantillon de taille n , il faut considérer que F_n joue le rôle de F et que F_n^* (la distrib. bootstrap) joue le rôle de F_n dans le processus de rééchantillonnage.

- L'échantillon de base ^{est composé de} n observ. iid de distrib. F . Notez que l'estimation du paramètre d'intérêt vaut ~~donc~~ $T(F_n)$. Donc en bootstrapping, on laisse F_n jouer le rôle de F et on prend n observations iid provenant de F_n - puisque F_n est la FdR empirique, il s'agit simplement de tirer aléatoirement avec remise des observations de l'échantillon d'origine. Car $P(X_i = x_i) = \frac{1}{n}$ à chaque tirage doit être satisfait. (d'où la remise).

- Ex: $n=5$, $\begin{matrix} X_1=7 & X_3=3 & X_5=6 \\ X_2=5 & X_4=9 \end{matrix} \Rightarrow \bar{x}_n = 6$ dans cet échantillon

Tirer aléatoirement avec remise dans cet échantillon génère un "échantillon bootstrap".

- ~~Exemple~~ Echantillon bootstrap: on le note $X_1^*, X_2^*, X_3^*, X_4^*, X_5^*$; cet échantillon permet de définir la distribution bootstrap de rééchantillonnage avec remise, notée F_n^* .

$\Rightarrow$ estimation bootstrap = $T(F_n^*)$.

Ex avec l'échantillon d'origine précédent: $\begin{matrix} X_1^* = 5 & X_3^* = 7 & X_5^* = 5 \\ X_2^* = 9 & X_4^* = 7 \end{matrix} \Rightarrow \bar{x}_n^* = 6,6$

Donc l'estimateur bootstrap = $6,6 \neq 6$ (échantillon origine)

- Remarquez que certaines valeurs/observ. de l'échantillon d'origine sont répétées ou omises dans l'échantillon bootstrap.
- Si on répète la génération d'un nouvel échantillon, bootstrap, on aura (probablement) une estimation de la moyenne encore différente.

- Si on répète l'opération plein de fois, on aura plein de valeurs $\neq$ pour $\bar{T}_n^*$. (3)
 $\Rightarrow$ on pourra faire un histogramme des valeurs de la moyenne
 $\Rightarrow$ on l'appelle l'approximation Monte Carlo de la distribution bootstrap.
 $\Rightarrow$ la moyenne de toutes ces valeurs (donc la moyenne des moyennes $\bar{T}_n^*$) sera proche de $\bar{T}$ puisque la moyenne théorique de la distrib. bootstrap est la moyenne de l'échantillon d'origine. Cependant, l'histogramme issu du rééchantillonnage permet de rendre compte de la variabilité de ces estimations, de leur asymétrie, écart-type, etc...

- En théorie, l'estimation bootstrap exacte du paramètre peut être calculée en moyennant sur tous les échantillons bootstrap possibles. (dans le cas précédent, l'estimation ^{exacte} vaudrait $\bar{T}$).
 Mais en pratique, puisqu'il y a n^n échantillons bootstraps distincts possibles, ce devient vite infaisable (ex: $n=10, \Rightarrow 10\,000\,000\,000!$) $\Rightarrow$ en pratique, on utilise une approx. Monte-Carlo!

- Si on génère aléatoirement M ($M=10\,000$ par ex.) échantillons bootstrap, la distrib. des estimations bootstrap approche la distribution bootstrap de l'estimation.
 Plus M est grand, plus l'histogramme approche la distribution bootstrap vraie de l'estimation.

- Fonctionnement de l'approx. Monte Carlo:
 ① Générer un échantillon par tirage avec remise depuis la distribution empirique des données, on a ^{ainsi} créé un échantillon bootstrap.
 ② Calculer $T(\bar{F}_n^*)$, l'estimation bootstrap de $T(F)$: on remplace l'échantillon d'origine par l'échantillon bootstrap et l'estimation bootstrap de $T(F)$ à la place de l'estimation de $T(F)$ via les données d'origine.
 ③ Répéter ① et ② M fois avec M grand.

⚠: IMPORTANT: l'application de Monte Carlo approx au bootstrap implique 2 sources d'erreur:

- 1- l'approx. Monte Carlo de la distribution bootstrap, qui peut être considérablement réduite si M est grand.
- 2- l'approximation de la distribution bootstrap $\bar{F}_n^*$ pour la distrib. originelle F de la population.

* Fin du
1er cours
de 2h

- Si $T(F_n^*) \xrightarrow{n \rightarrow \infty} T(F)$ alors le bootstrapping fonctionne.
 $\Rightarrow$ ça marche dans bc de cas, mais pas toujours. $\mathbb{P}(\lim_{n \rightarrow \infty} \|F_n - F\|_\infty = 0) = 1$
- Rappel: Th. de Glivenko-Cantelli: $F_n \rightarrow F$ uniformément: $(\sup_x |F_n(x) - F(x)| \xrightarrow{n \rightarrow \infty} 0?)$
- Souvent nous savons que l'estimateur empirique est consistant. Par ex.
 $\frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{n \rightarrow \infty} E(X)$ d'après la loi des grands nombres.

Donc on a (1) $T(F_n) \xrightarrow{n \rightarrow \infty} T(F)$.

Mais cela dépend aussi de la fonctionnelle T et de sa régularité (dérivée, ...)

$\Rightarrow$ On a donc aussi besoin de (2) $T(F_n^*) \xrightarrow{n \rightarrow \infty} T(F_n)$, plus précisément

$$(2) \quad T(F_n^*) - T(F_n) \xrightarrow{n \rightarrow \infty} 0$$

⇒ Rq: Le bootstrap reste basé sur une théorie asymptotique: des études simulatoires où la taille de l'échantillon est petite sont nécessaires pour évaluer la pertinence de son usage.

① - Le Jackknife

- Introduit par Quenouille (1948)
- Le but était d'améliorer l'estimation en corrigeant ses biais. [estimateur $\rightarrow$ biais $\rightarrow$ variance élevée]
- Tukey (1958) montre ensuite que le jackknife peut aussi être utilisé pour estimer l'écart-type d'une estimation.
- Le nom "jackknife" (couteau suisse) vient du fait que la technique peut être utilisée pour $\neq$ objectifs.
- Le bootstrapping utilise des échantillons bootstrap pour estimer une variabilité de l'estimation, alors que le jackknife utilise les "pseudo-observations".

Question: "l'estimateur bootstrap" est-il $T(F_n^*)$? Non

ou la moyenne des $T(F_n^*)$? oui (cf ex. application du Monte Carlo)

• Tout d'abord, notons $\tilde{\mu}$ une estimation basée sur un échantillon d'observations iid de taille n , $(X_1, \dots, X_n)$ où $X_i \sim F$.

Comme précédemment, on note F_n la FdR empirique et

$$\begin{cases} \mu = T(F) \text{ une fonctionnelle} \\ \tilde{\mu} = T(F_n) \\ \tilde{\mu}_{(i)} = T(F_{n(i)}) \text{, où } F_{n(i)} \text{ est la FdR empirique pour l'échantillon sans l'observ. } i. \end{cases}$$

ex: $\tilde{\mu}$ est la variance de la population (de l'échantillon), l'estimation jackknife de variance de σ^2 vaut ($\sigma^2 = \text{Var}_{\text{population}}(\tilde{\mu})$)

→ variance de la variance: $\sigma_{\text{JACK}}^2 = \frac{n}{n-1} \sum_{i=1}^n (\tilde{\mu}_{(i)} - \mu^*)^2$ avec $\mu^* = \frac{1}{n} \sum_{i=1}^n \tilde{\mu}_{(i)}$

l'estimation jackknife de l'écart-type pour $\tilde{\mu}$ vaut $\sqrt{\sigma_{\text{JACK}}^2}$.

la variance
variance estimée sans l'observ. i.
moyenne des variances

• Tukey définit les pseudo-observations, ou pseudo-valeurs:

en fait générale

$$\tilde{\mu}_i = \tilde{\mu} + (n-1)(\tilde{\mu} - \tilde{\mu}_{(i)})$$

pour revenir au même que précédemment mais aboutir à une erreur type Monte Carlo

⇒ l'estimation jackknife du paramètre μ vaut $\mu_{\text{JACK}} = \frac{1}{n} \sum_{i=1}^n \tilde{\mu}_i$ (même principe que Monte Carlo)

variance empirique

l'estimation est la moyenne des pseudo-valeurs.

• On écrit par ex l'estimation de la variance σ^2 de l'estimation $\tilde{\mu}$ avec les pseudo-valeurs:

$$\sigma_{\text{JACK}}^2 = \frac{1}{n(n-1)} \sum_{i=1}^n (\tilde{\mu}_i - \mu_{\text{JACK}})^2$$

moyenne
variance de la variance
moyenne des variances
= moyenne des variances de variance

⇒ dans cette forme, on voit que la variance est l'estimateur usuel pour la variance d'une moyenne. Dans ce cas, c'est la moyenne des pseudo-valeurs.

• le lien entre l'estimation jackknife et l'estimateur de l'écart-type est une estimation bootstrap où l'estimation du paramètre d'intérêt est remplacé par une approx. linéaire de lui-même. ⇒ il y a donc un lien étroit entre les 2...

② - la méthode Delta: (estimateur fonction d'un autre estimateur)

On est souvent intéressé par les moments d'un estimateur. En particulier, c'est souvent la variance de l'estimateur qui nous intéresse.

Pour illustration, soit $\varphi = f(\alpha)$ avec $(\alpha, \varphi) \in \mathbb{R}$
 f une fonction dérivable par rapport à α .

$\Rightarrow$ $\exists$ une expansion de Taylor en α_0 : au 1^{er} ordre cela donne, pour $\alpha \in \mathcal{N}(\alpha_0)$,

$$\varphi = f(\alpha) = f(\alpha_0) + (\alpha - \alpha_0) f'(\alpha_0) + o(\alpha - \alpha_0)$$

$$\Rightarrow \varphi = f(\alpha) = f(\alpha_0) + (\alpha - \alpha_0) f'(\alpha_0)$$

$$\text{ou } f(\alpha) - f(\alpha_0) = (\alpha - \alpha_0) f'(\alpha_0)$$

$$\Rightarrow (f(\alpha) - f(\alpha_0))^2 = (\alpha - \alpha_0)^2 (f'(\alpha_0))^2$$

Mainenant, disons que φ est une v.a., on prend l'espérance dans l'éqval. précédente:

$$E[(f(\alpha) - f(\alpha_0))^2] = E[(\alpha - \alpha_0)^2] \underbrace{(f'(\alpha_0))^2}_{\text{cte}}$$

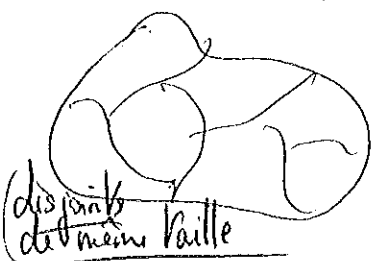
Cette dernière équation est l'approximation par la méthode Delta de la variance de $\varphi = f(\alpha)$ puisque le terme de gauche est à peu près la variance de φ et celui de droite la variance d' α multipliée par la constante $(f'(\alpha_0))^2$ si l'on choisit α_0 comme la moyenne de α .

③ - La validation croisée

La validation croisée est une procédure générale en modélisation statistique.

Elle est ~~par exemple~~ utilisée pour de la sélection de modèle parmi un ensemble de modèles, par exemple pour:

- ordre d'un processus AR
- select. de variable en régression logistique / régression linéaire multiple
- nombre de composantes d'un modèle mélange fini
- choix d'un modèle de classif. paramétrique
- élagage d'arbre de classification, ...



L'idée est de partager aléatoirement les données en 2 sous-ensembles: 1 pour (calibrer / entraîner) le modèle, l'autre pour le valider. On moyenne ensuite les erreurs du modèle données sur toutes ces combinaisons pour choisir le meilleur modèle.

④ le sous-échantillonnage

⑤

- Développé par Hartigan (1968) $\Rightarrow$ Théorie des intervalles de confiance par sous-échantillonnage aléatoire.
- Récemment, cette théorie a été reliée au bootstrap, pour des données indépendantes ou dépendantes. Voir Politis & al. (1999); Politis and Romano (1994), ...
- Définition: soient $S_1, S_2, \dots, S_{B-1}$; $(B-1)$ des $(2^n - 1)$ sous ensembles non-vides des entiers $1, 2, \dots, n$.



$\times 3$

1, 2 - 3	1, 2, 3	2
1, 3 - 2	1, 2, 3	3
2, 3 - 1	1	

$2^3 - 1 = 7$ possibilités

Ces $(B-1)$ sous-ensembles sont sélectionnés aléatoirement sans remise, et les estimations sont ensuite faites.

Preuve:

~~Théorème des combinaisons~~

$E = \{1, 2, \dots, n\}$

Si $\text{Card}(E) = n$ (ce qui est notre cas dans notre exemple)

alors $\text{Card}(P(E)) = 2^n$ où $P(E)$ est l'ensemble des parties de E .

Par récurrence: $n=0 \Rightarrow$ par d'élément $\rightarrow$ l'ensemble-vidé a un seul sous-ensemble (lui-même).
 $\Rightarrow 2^0 - 1$ donc OK

- supposons la prop. vraie pour $n=m \Rightarrow 2^m$ parties
- pour $n=m+1$: l'ensemble E est donc non-vidé. Soit a un élément de E .

ou preuve: on peut aussi compter les parties à 0, 1, 2, ..., n éléments d'un ensemble E à n éléments:

$$\text{Card}(P(E)) = \underbrace{\binom{0}{n}}_{\text{ensemble-vidé}} + \binom{1}{n} + \binom{2}{n} + \dots + \binom{n}{n} = (1+1)^n = 2^n$$

Si l'on enlève l'ensemble-vidé, on a donc $2^n - 1$ parties.

Puis on poursuit le même raisonnement que précédemment à partir des échantillons créés?

III - Variété des applications

- On est tenté d'appliquer le bootstrap un peu partout... mais on a vu que cela ne marche pas toujours... alors comment le sait-on ? $\rightarrow$ soit prouvé théoriquement (théorème de CV) soit prouvé par simus.

Par exemple en régression, il y a 2 approches au bootstrapping :

- 1) "bootstrapping résiduels" : on fit le modèle puis on bootstrapte les résidus.
- 2) "bootstrapping vectors or cases" : on bootstrap sur $(Y_i, X_{i1}, \dots, X_{iR_i})$ pour $i=1, \dots, n$ puis on fit chaque fois un modèle $\Rightarrow$ mieux si on a des doutes sur la forme paramétrique choisie pour la modélisation.

- Quelque soit l'allure des données, le bootstrap peut fonctionner : ^{in en} cas theory-violated, highly skewed, ... puisqu'il n'y a pas d'hyp-param.
- Simplicité de mise en œuvre : Monte Carlo... et ordinateurs puissants aujourd'hui.
Mais Δ : il n'est pas toujours facile de voir quand le bootstrap ne fonctionne pas, et de savoir pourquoi...

IV - Bootstrap et R

- Beaucoup de librairies qui traitent du sujet : $\left(\begin{array}{l} > \text{help.search("bootstrap")} \\ ?? \text{bootstrap} \end{array} \right)$
- Package "bootstrap"
"boot"
- Pour le rééchantillonnage, la fonction de base est "sample(.)".

V - Histoire et biblio

cf p 25

VI - Exercices

cf p. 26

Exemple: considérons l'estimateur du max. de vraisemblance de la variance σ^2 d'une n -e. gaussienne: $X \sim \mathcal{N}(\mu, \sigma^2)$

$$S_n^2 = \frac{1}{n} \sum_{i=1}^n (X_i - X_b)^2 \quad \text{avec} \quad X_b = \frac{1}{n} \sum_{i=1}^n X_i$$

Cet estimateur a un biais $b = -\frac{\sigma^2}{n}$ connu ici et dont l'expression analytique existe!

→ la (vraie) estimation bootstrap du biais vaut b (comme concentration = 6 tout à l'heure).

→ On définit $B^* = E(\theta^* - \hat{\theta})$ avec $\begin{cases} \hat{\theta} = S_n^2 \text{ estimateur du max de vraisemblance} \\ \theta^* = \frac{1}{n} \sum_{i=1}^n (X_i^* - X_b)^2 : \text{estimateur bootstrap} \end{cases}$

L'approx. Monte Carlo de B^* vaut $B_{MC} = \frac{1}{N} \sum_{j=1}^N B_j^*$ avec B_j^* une estimation du biais pour le j^e échantillon bootstrap.

Ainsi $B_j^* = \theta_j^* - \hat{\theta}$ et $\theta_j^* = \frac{1}{n} \sum_{i=1}^n (X_{ij}^* - X_b)^2$.
 X_{ij}^* est l' i^e observation du j^e échantillon bootstrap. moyenne de l'échantillon d'origine (11: pas celle de l'échantillon bootstrap!)

Illustration 1: machine

- Le but de détecter un biais est de corriger un estimateur en le débiaisant.
- 1. la réduction du biais se fait en général au détriment de l'↑ de la variance de l'estimateur.
 Il faut débiaiser l'estimateur si le carré du biais est $>$ ↑ de la variance $\Rightarrow$ MSE de $\hat{\theta} \downarrow$
 sinon ce n'est pas forcément pertinent. $\downarrow$
car $MSE \approx \text{biais}^2 + \text{variance}$

② Exemples en analyse linéaire discriminante (classification error rate)

③ Exemple simple en analyse linéaire discriminante et estimation du taux d'erreur bootstrap

On imagine 2 classes dans la population (de 10 individus):

→ 5 individus $\in$ class 1, où la class 1 $\sim \mathcal{N}((0,0)^T, \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix})$

→ 5 " " " 2, " " $\sim \mathcal{N}((1,1)^T, \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix})$

On génère les individus - pour chaque classe (on connaît donc ici les classes!):

→ class 1: $(2,052; 0,339)^T; (1,083; -1,32)^T; (0,083; -1,524)^T; (1,278; -0,452)^T;$

→ class 2: $(1,307; 2,268)^T; (-0,548; 1,741)^T; (2,498; 0,813)^T; (0,832; 1,409)^T; (-1,226; -0,606)^T;$
 $(1,498; 2,063)^T$

⇒ On va générer 6 échantillons bootstrap de taille 10 (taille de la population), et

Exercice 1: on retrouve dans le beurre de cacahuète des résidus d'Aflatoxine. Dans les 12 lots de beurre de cacahuètes à disposition, on mesure la concentration d'Aflatoxine résiduelle:

4,34	5,06	4,53	5,07	4,38	5,16	4,38	4,43
4,33	4,72	4,32	4,36				

(6)

a) Combien d'échantillons bootstrap possibles peut-on créer?

R: 12^{12} puisqu'on en tire 12 et qu'il y a 12 possibilités à chaque tirage.

b) En utilisant R et fonction `sample()`, générer 5 échantillons d'entiers entre 1 et 12. de taille 12.

c) Pour chacun des échantillons, créés, calculer la moyenne de la concentration en aflatoxine.

d) Calculer la moyenne de ces moyennes, et comparer la à la moyenne de l'échantillon d'origine.

e) Trouver le minimum et le maximum des 5 moyennes provenant des échantillons bootstrap.

C'est un intervalle de confiance grossier sur la moyenne - Si vous avez créé 1000 échantillons bootstrap, un intervalle de confiance à 95% de la moyenne aurait été défini par les 25^e et 975^e moyennes lorsqu'elles sont classées par ordre croissant.

2. ESTIMATION

On parlera ici d'estimation ponctuelle, et de comment le bootstrap peut être utilisé pour estimer le biais d'une estimation, afin de l'améliorer après ajustement.

Historiquement, le bootstrap a été introduit pour estimer l'écart-type d'une estimation, puis pour l'ajustement d'un biais (contrairement au jackknife).

I. Estimation du biais

① Ajustement bootstrap

$$E[\underbrace{\hat{\theta}}_{\text{v.a.}} - \underbrace{\theta}_{\text{cste}}] = 0 \Rightarrow \hat{\theta} \text{ estimateur sans biais de } \theta.$$

Soit $E(X)$ l'espérance de X . Soit $\hat{\theta}$ une estimation du paramètre θ .

$$\text{Posons } X = \hat{\theta} - \theta$$

Le biais d'un estimateur $\hat{\theta}$ de θ est donné par $b = E(X) = E(\hat{\theta} - \theta)$.

Rp: la distrib bootstrap d'un estimateur peut donner une idée de la loi de l'estimateur si la bootstrap s'applique.
(ex: EMV qui est gaussien se retrouve).

II - Estimer un paramètre de localisation

Cette partie parle de l'estimation ponctuelle d'indicateurs de localisation (moyenne, médiane).
Dans le cas général de distributions "classiques" (moments finis jusqu'à l'ordre 4), la moyenne est le meilleur estimateur de localisation, et le plus adopté pour les distrib symétriques.
Mais cet indicateur de localisation n'est pas toujours défini (ex: loi de Cauchy), et n'est pas le meilleur indicateur si les distrib présentent de fortes asymétries.

① Estimer une moyenne:

- Pour les populations avec un moment d'ordre 1 fini, on peut estimer la moyenne.
 - Pour des distrib. symétriques, la moyenne est la meilleure mesure pour la tendance centrale.
- Comment le bootstrap se comporte-t-il pour déterminer cette estimation?

- Dans le cas gaussien par ex., la moyenne empirique est l'EMV de μ ($X \sim \mathcal{N}(\mu, \sigma^2)$).
La moyenne empirique est aussi l'ESBM (Estimateur Sans Biais et la Variance Minimale) de μ .
Comment le bootstrap estimateur peut-il faire mieux? $\rightarrow$ en fait il ne peut PAS!

Dans ce cas précis, on peut déterminer l'estimation bootstrap sans approx. Monte Carlo.
En effet car la moyenne empirique est aussi la moyenne de toutes les estimations bootstrap de moyennes empiriques $\rightarrow$ l'estimation bootstrap vaut la moyenne empirique!

② - Estimer une médiane

Dans le cas de distributions fortement asymétriques, la médiane ou le mode sont de meilleurs indicateurs de tendance centrale.

Par exemple pour la distribution de Cauchy, la médiane ^{théorique} ~~empirique~~ est bien définie et la médiane empirique ("sample median") est un estimateur consistant de la médiane théo. ("population median").

Dans ce cas aussi, l'estimateur bootstrap de la médiane est aussi ^{CV} consistant, mais et n'a pas d'avantages supplémentaires.

Cependant, si l'on veut estimer l'écart-type des indicateurs (moyenne, médiane), le bootstrap permet de le quantifier.

III - Estimer la dispersion

de la distribution de la population? ou de la distrib. de l'estimateur?

Quand le moment d'ordre 2 existe, l'écart-type peut être estimé.

Pour des estimateurs non-biaisés, cet écart-type est appelé "standard error of the estimate".

L'écart-type est souvent considéré comme le meilleur indicateur de dispersion, en tout cas pour des distributions classiques. De plus, cet indicateur intervient dans des résultats fondamentaux (ex: inégalité de Chebyshev) qui nous aide à comprendre la dispersion des données.

Malheureusement, cet indicateur n'existe parfois pas (ex: Cauchy). On se base alors sur un autre indicateur: l'"interquartile range".

① Estimer l'écart-type d'un estimateur:

- Certains auteurs ont donné un calcul exact de l'estimateur bootstrap de l'écart-type d'une médiane par exemple. Il suffit ensuite de le comparer avec l'approximation Monte Carlo du bootstrap sur un échantillon réel.

Application: pour simplifier, considérons un nombre impair d'observations X_i : $n = 2m + 1$

Pour estimer la médiane, ce sera donc $X_{(m+1)}$, où $X_{(i)}$ désigne la i^{e} observation classée par ordre croissant.

On ~~peut~~ donne la variance exacte de la distribution bootstrap de la médiane en utilisant des propriétés célèbres de statistiques d'ordre. (8)

Soit $X_{(1)}^*, X_{(2)}^*, \dots, X_{(n)}^*$ la statistique d'ordre d'un échantillon bootstrap issu de $X_{(1)}, \dots, X_{(n)}$.

Soit $X_{(i)}$ la i^e plus petite observation de l'échantillon d'origine, et

$$N_i^* = \# \{ j : X_j^* = x_{(i)}, j = 1, 2, \dots, n \}.$$

$\Rightarrow N_i^*$ est le nombre d'observations dans l'échantillon bootstrap égales à l'observation i ordonnée dans l'échantillon d'origine.

On peut montrer que $\sum_{i=1}^n N_i^* \sim \text{Binomiale}(n, p = \frac{k}{n})$.

Maintenant, soit p^* la probabilité sous rééchantillonnage bootstrap - On a :

$$p^* \left(\underbrace{X_{(m+1)}^*}_{\substack{\text{médiane} \\ \text{échantillon bootstrap}}} > x_k \right) = p^* \left(\underbrace{\sum_{i=1}^k N_i^*}_{\sim \text{Binomiale}} \leq n \right) = \sum_{j=0}^n C_n^j \left(\frac{k}{n} \right)^j \left(\frac{n-k}{n} \right)^{n-j}$$

En utilisant des relations entre sommes binomiales et fonctions bêta incomplètes, on peut déterminer une formule explicite pour l'estimateur bootstrap de l'écart-type de la médiane.
 $\Rightarrow$ On n'a pas besoin de approx. Monte Carlo.

• Dans d'autres problèmes d'estimation de paramètres, on n'a pas de formule explicite. L'approx. Monte Carlo est alors nécessaire.

(Soit $\hat{\theta}$ l'estimateur (basé sur les observ.) de θ .)

(Soit $\hat{\theta}_i^*$ l'estimateur bootstrap de θ pour l'échantillon bootstrap i .)

(Soit $\bar{\theta}^*$ la moyenne des $\hat{\theta}_i^*$.)

Alors, comme décrit dans Efron (1982), l'estimateur bootstrap de l'écart-type de l'estimateur

$\hat{\theta}$ est donné par

$$SD_b = \sqrt{\frac{1}{B-1} \sum_{k=1}^B (\hat{\theta}_k^* - \bar{\theta}^*)^2}$$

B : nb échantillons bootstrap.

issu de la distrib. bootstrap de l'estimateur.

$\hookrightarrow$ écart-type de l'estimateur bootstrap de l'estimation

② Estimer "l'interquartile range"

Un moyen naturel d'estimer cette quantité est de prendre le 25^e et le 75^e percentiles de la distribution `bookshop` pour obtenir l'estimation. Il faut donc au préalable classer les estimations `bookshop`. Application : machine.

IV - Régression linéaire