

Méthodologie, recueil et manipulation de données

Christophe Pallier

September 7-8, 2017

Preliminaires

Plan

1. Une peu d'épistémologie (Théorie de la connaissance)
2. Définitions d'une population et stratégies d'échantillonnage
3. Construction de questionnaires
4. Manipulation et exploration de données

Ressources



Partie 1: Un peu d'épistémologie

La Méthode scientifique

- Une bonne théorie doit faire des **prédictions testables**.

“Si les pierres qu’on trouve actuellement sur Terre tombent quand on les lâche, c’est parce que celles qui ne tombaient pas se sont toutes envolées dans l’espace.”

Rien de plus pratique qu’une bonne théorie !

La théorie dit **quoi regarder**.

Si vous n’avez pas théorie du fonctionnement d’un téléviseur, il est difficile de le dépanner (Métaphore du “dépannage”. voir aussi: développement logiciel, médecine, etc.)

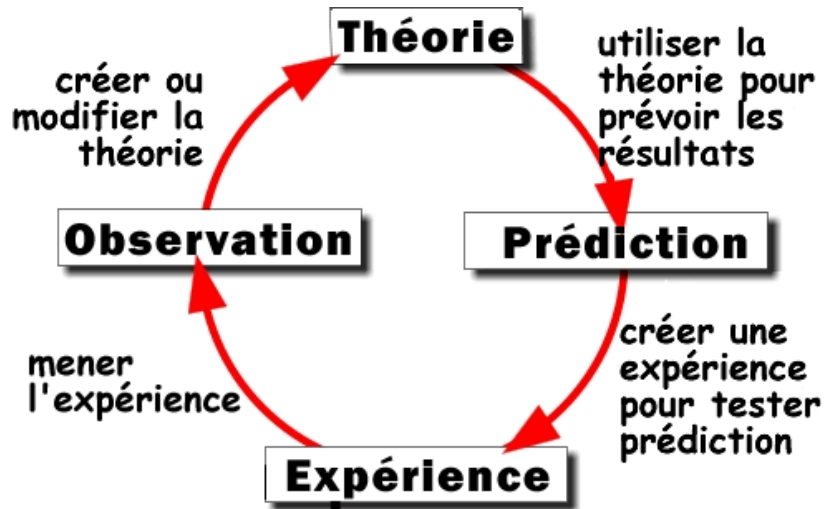


Figure 1: Des allers-retours entre Observation/Expérimentation et Théorie

Cependant, au début d'une recherche (dans sa *phase exploratoire*), il existe souvent des questions *a- théoriques* (ou, plus positivement, des questions *empiriques*).

Questions "a-théoriques"

Des questions purement *empiriques*, ou *descriptives*, peuvent présenter de l'intérêt (ou non). Exemples:

- L'espérance de vie dépend-elle du niveau socio-économique?
(il faut le mettre en évidence avant de rechercher les causes: meilleure hygiène de vie, soins, génétique,...)
- Y-a-t il un lien entre l'âge d'acquisition d'une seconde langue et le niveau qu'on peut atteindre (indépendamment des autres facteurs: quantité d'exposition à l2, type d'enseignement) ?
- Un champ magnétique peut-il avoir des effets thérapeutiques ("magnétothérapie") ?
- Les champs électro-magnétiques ont-ils des effets délétaires sur la santé ?
- La taille ou la couleur de cheveux d'un individu ont-ils des effets sur son revenu ?

Observation vs. Expérimentation

Observer: Noter et rapporter les événements de manière systématique.

Expérimenter: *manipuler des conditions* et démontrer que des événements se reproduisent sous certaines d'entre elles.

Exemples de variables ne pouvant pas être manipulées directement:

- la taille d'une étoile et sa couleur.
- le poids, la taille, le sexe d'un individu.
- Le fait d'avoir une certaine maladie psychiatrique ou non.

Exemples de variables pouvant être manipulées (assignées arbitrairement à des individus):

- médicament actif vs. placebo
- aide sociale (économie expérimentale). Voir <https://80000hours.org/articles/effective-social-program/>

Corrélation et causalité

- Il est bien connu que la mise en évidence d'une corrélation n'implique pas de lien de causalité directe.

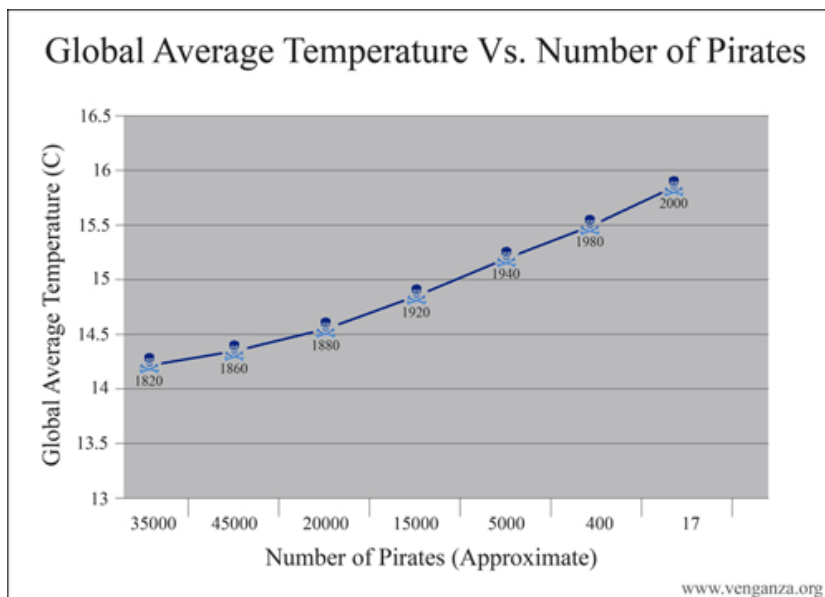
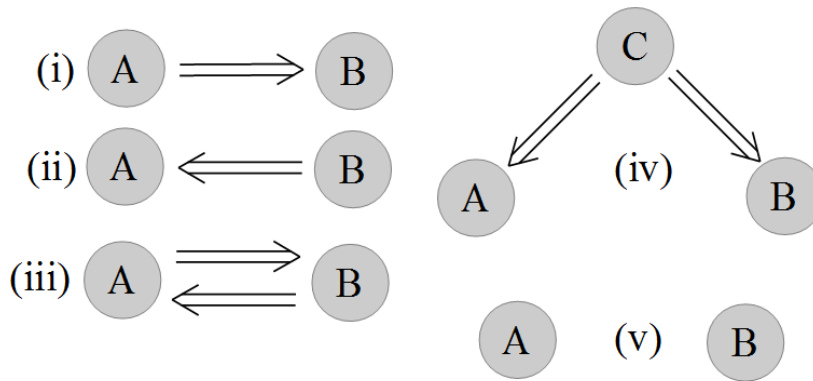


Figure 2: effet du nombre de Pirates sur le réchauffement climatique

- Par exemple, s'il existe une corrélation entre la cylindrée de la voiture familiale et le QI des enfants, cela est certainement dû à des variables tierces (niveau socio-économique, ... ; on s'en doute uniquement parce qu'on a une théorie!)

Figure 3: relations entre variables



Causalité et corrélation : quelques possibilités

- Seul l'expérimentation permet de tester, en terme de causalité, l'impact d'une ou plusieurs variable(s) indépendantes sur une ou plusieurs variables dépendantes.
- Néanmoins, même avec des données d'observation, on peut calculer des *coefficients de corrélations partiels* et ajuster pour les effets de variables de "contrôles".

Claude Bernard (1877) Principes de médecine expérimentale.{smaller}

"Dans toutes les sciences, il y a deux états bien distincts à considérer.

Ce sont :

1. l'état de science d'observation ;
2. l'état de science expérimentale.

Ces deux états sont nécessairement et absolument subordonnés l'un à l'autre. Jamais une science ne peut parvenir à l'état de science expérimentale sans avoir passé par l'état de science d'observation. Mais il y a des sciences auxquelles il n'est pas donné de pouvoir parvenir à l'état de science expérimentale ; telle est l'astronomie, par exemple.

Une science d'observation, ou science naturelle, se borne à observer, à classer, à contempler les phénomènes de la nature et à déduire des observations les lois générales des phénomènes. Mais elle n'agit pas sur les phénomènes eux-mêmes pour les modifier ou en créer de nouveaux, pour agir sur la nature en un mot. La science d'observation est une science passive ; elle prévoit, se gare, évite, mais ne change rien activement. Or, les sciences qui, comme

l'astronomie, s'occupent de phénomènes hors de notre portée expérimentale, restent forcément des sciences d'observation.

Les sciences expérimentales, au contraire, sont plus ambitieuses ; elles veulent agir et étendre leur puissance sur la nature, modifier les phénomènes, en créer qui n'existent pas et régler les éléments à leur volonté.

Conséquence: Claude Bernard ne prenait pas au sérieux la théorie de l'évolution de Darwin, ce qui peut être un peu trop extrémiste.

Conséquences pratiques

Choix de la procédure statistique pour quantifier la degré de relation linéaire entre 2 variables :

- Dans une situation d'observation, on privilégie la *corrélation simple* où les deux variables jouent un rôle symétrique.

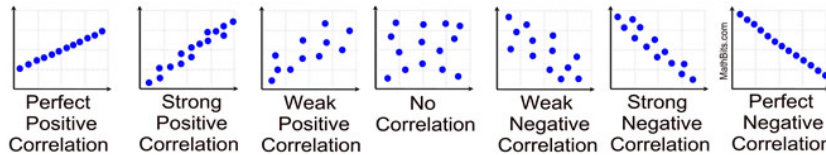


Figure 4: Corrélation linéaire entre deux variables

On calcule typiquement le coefficient de corrélation de Pearson :

$$r(X, Y) = \frac{\sigma_{X,Y}^2}{\sigma_X \cdot \sigma_Y}$$

- Dans une situation expérimentale où les variables jouent des rôles *asymétriques* — avec une variable “dépendante” et une variable “indépendante” —, on utilisera plutôt la *régression*.

$$Y_i = aX_i + b + \epsilon_i$$

Remarques sur le coefficient de corrélation

Le coeff. de corrélation ne reflète *que* la composante linéaire. Même quand il est nul, il peut néanmoins exister une *relation non-linéaire* entre les deux variables:

Ne jamais rapporter un coefficient de corrélation sans examen graphique !

Pour en savoir plus : https://en.wikipedia.org/wiki/Correlation_and_dependence

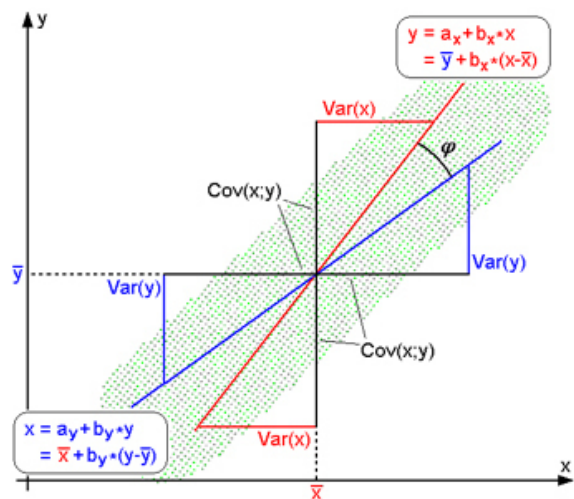


Figure 5: Les droites de régression de X sur Y et de Y sur X diffèrent

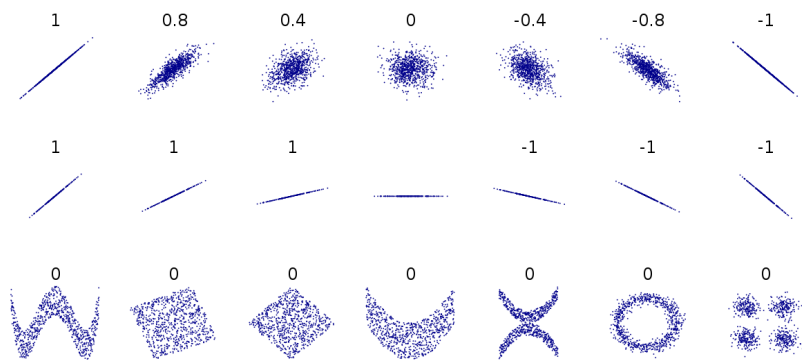


Figure 6: Exemples de relations possibles entre deux variables continues

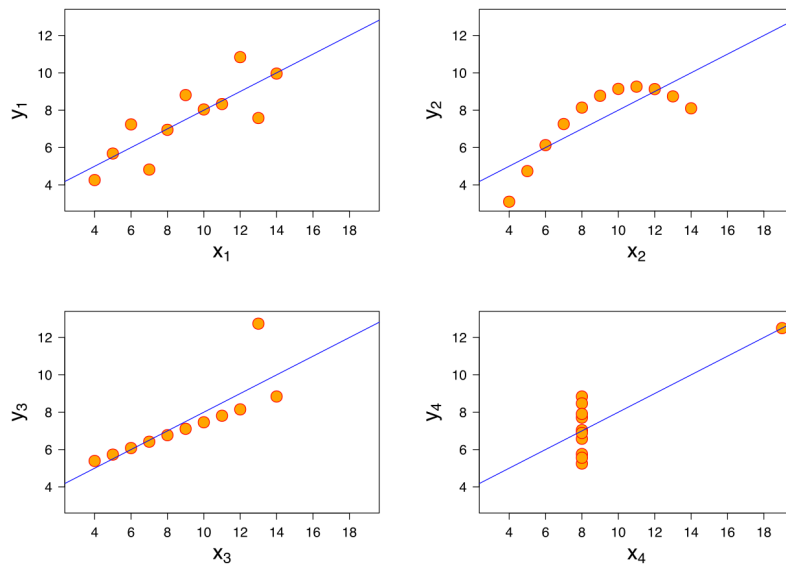


Figure 7: Les coefficients de corrélation sont exactement les mêmes dans ces quatre figures

Avantages de l'observation sur l'expérimentation

Selon notre définition, les enquêtes ne font pas partie des recherches expérimentales: on n'y manipule pas de variables, on se contente de les enregistrer.

Néanmoins, les enquêtes présentent certains avantages sur les travaux expérimentaux:

- Les enquêtes sont typiquement moins coûteuses par individu, ce qui permet d'obtenir des échantillons plus larges et plus représentatifs (Les expériences sont typiquement réalisées sur des étudiants sains et volontaires)
- Le résultat d'une enquête (bien menée) renseigne sur le "monde réel" alors que la généralisabilité des expériences en laboratoire peut être mise en doute (manque de validité écologique).

En pratique, les deux approches sont complémentaires. Toutes les méthodes sont imparfaites, et des *résultats convergents* selon plusieurs méthodes sont très appréciables.

Démarche exploratoire vs. descriptive, explicative

Induction et Déduction

- **Le scandale de l'induction:** aller du particulier au général n'est pas *logiquement* valide.

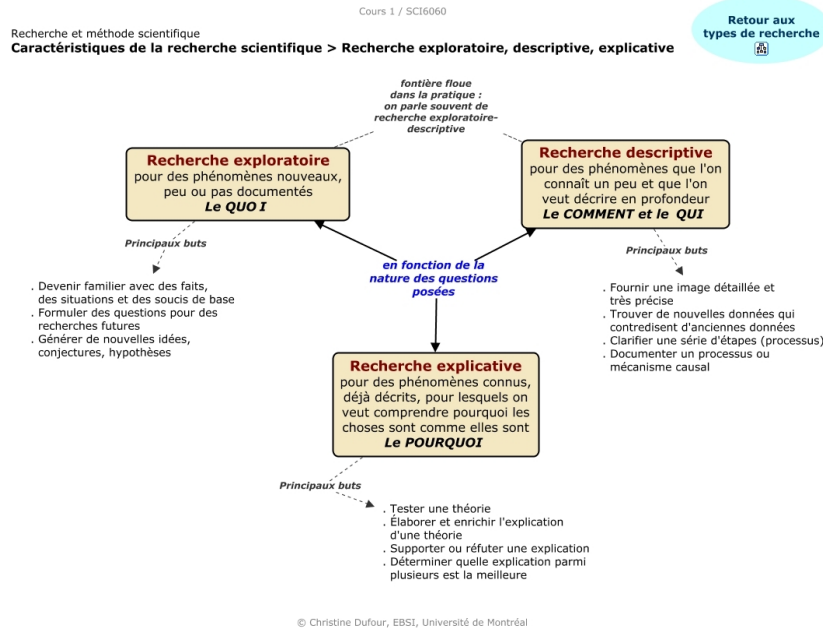


Figure 8: Il est important de bien identifier le type de démarche

Un physicien, un mathématicien et un logicien voyagent dans un train qui traverse la campagne. Ils croisent un champ où se trouve une vache blanche.

Le physicien: "tiens, dans ce pays, les vaches sont blanches!" Le mathématicien: "Non: il y a au moins une vache blanche dans ce pays." Le logicien: "Je te corrige: il y a au moins une vache avec un côté blanc."

- Seule la **déduction** est logiquement valide.

Tous les hommes sont mortels, Socrate est un homme, donc Socrate est mortel.

Et pourtant vive l'induction !

En réalité, on peut souvent faire le pari qu'il y a des régularités dans les événements:

"Jusqu'à présent, le soleil s'est levé tous les matins, donc il se lèvera tous les jours (ou au moins avec une forte probabilité)."

Intuitivement, si on a observé un événement **A**, il y a plus de chance qu'il se reproduise qu'un événement **B** qui n'a jamais été observé:

En pratique, scientifiques et autres utilisent l'induction *en essayant de quantifier notre incertitude* à l'aide de méthodes statistiques (Analyse risques/bénéfice. cf. Assurances)

En statistiques *fréquentistes*, la décision d'accepter un effet comme significatif est prise avec un seuil de risque, typiquement 5%.

En statistiques *Bayésienne*, on quantifie l'(in)certitude de notre connaissance sur des paramètres.

Karl Popper et la Falsification

Pour Karl Popper, une théorie scientifique doit pouvoir être réfutée par une observation (un "test crucial"). L'avancée des connaissances se fait en rejetant des théories, celles qui ne sont pas encore falsifiées sont les meilleures à un moment donné.

- L'astrologie, la psychanalyse, l'astronomie, l'histoire, ... font-elles des prédictions falsifiables ?

En réalité, les choses sont beaucoup plus complexes: On ne rejette pas une théorie dès qu'une expérience la contredit. On essaye d'abord de rejeter des hypothèses auxiliaires, c.à.d. de modifier le modèle (Paul Duhem).

Pour aller plus loin:

- Thomas Kuhn (1962) *La structure des Révolutions Scientifiques*
- Imre Lakatos (1976) *Conjectures et Refutations*
- Paul Feyerabend (1975) *Contre la méthode. Esquisse d'une théorie anarchiste de la connaissance*

Vie et mort des théories

- Il arrive que plusieurs théories soient compatibles avec les données. On cherche alors si elles font des prédictions différentes, et on teste ces prédictions pour les départager. (Voir John R. Platt, *Strong Inference, Science*, 1964).
- Une théorie peut évoluer en l'absence de nouvelles données, mais sur des critères formels internes. P.ex. en linguistique, on choisit entre plusieurs modèles du langage en fonction de leur apprenabilité.

Objectivité scientifique

Nos sentiments personnels ou nos attentes ne doivent pas influencer les données que nous rapportons (p.ex. biais de sélection des données). (D'où l'importance des expériences en *double aveugle*). Les articles scientifiques distinguent la partie méthodes/résultats des parties interprétatives.

Lecture recommandée: Steven J. Gould *La mal-mesure de l'Homme*
Odile Jacob.



Figure 9: La pyramide de la qualité des preuves

La démarche scientifique en pratique

- Choisir une problématique
- Faire une revue de l'état de l'art (synthèse de la littérature scientifique)
- Suggérer une ou plusieurs hypothèses *opérationnelles* (une définition opérationnelle spécifie les procédures effectives de mesure d'une variable. Ex: 'anxiété')
- Construire un plan d'acquisition de données
- (déposer des demandes de financement; recommencer; recommencer; ...)
- Recueillir les données
- Tester le(s) hypothèses, en utilisant des outils d'analyse statistique.
- Rédiger un compte rendu dans un article ou rapport de recherche.
- Evaluation par les pairs
- Publication
- Quelquefois: Reproduction par d'autres chercheurs

Principales étapes d'une enquête

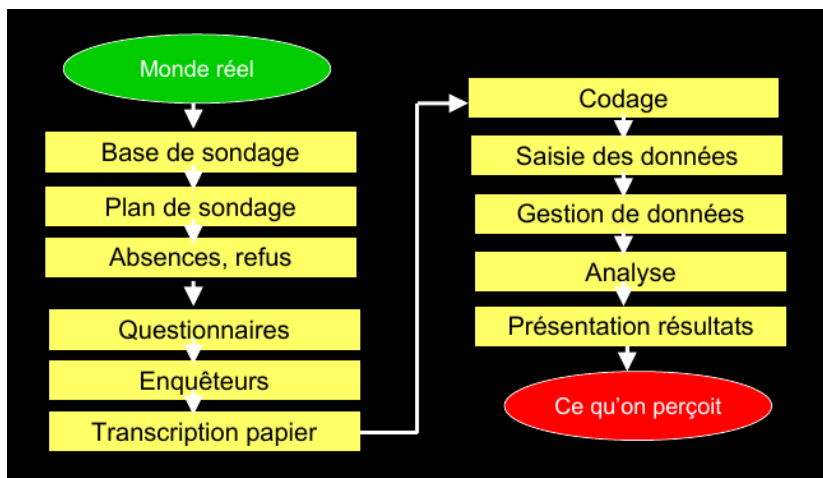


Figure 10: Principales étapes d'une enquête

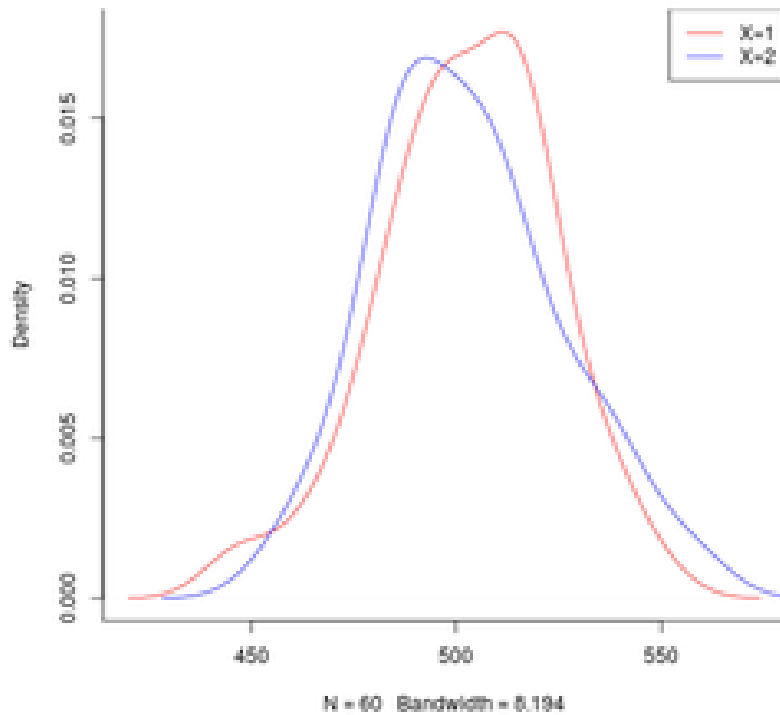
Partie 2: Définition d'une population et constitution d'un échantillon

Effet d'une variable

- Que signifie, d'un point de vue statistique, qu'une variable "à un effet sur une autre" ?

- Y est influencée par X si la *distribution de Y change quand X prend des valeurs différentes.

Figure 11: Distribution de Y en fonction de X



Population-mère et échantillon

On a rarement les moyens de contacter tous les membres de la population d'intérêt, appelée **population-mère** (sauf en cas de "recensement").

On doit se limiter à une population plus réduite (**l'échantillon*) qui est censée représenter la population-mère et qui doit nous permettre de *généraliser* — par induction — les résultats observés.

Problème: si pour comparer deux populations, on compare deux échantillons, il est quasiment certain qu'on va observer des différences entre les échantillons. *Mais celles-ci reflètent-elles des différences réelles entre les populations?*

Autrement dit, si on tire deux échantillons d'une même population mère, ils seront presque certainement différents à cause de la variabilité des individus.

Exemple: les Français et les Allemands ont-ils la même taille moyenne ? Si on prend 100 individus français et 100 individus allemands, il est peu probable que la taille moyenne soit égale au millimètre près.

Statistiques descriptives et statistiques inférentielle

Remarque: Il est très important de distinguer deux types d'usage des statistiques:

- **descriptives:** qui fournissent des indicateurs qui décrivent les données que l'on possède effectivement (moyenne, écart-type, ...)
- **inférentielles:** qui cherchent à inférer des propriétés des populations mères à partir des données observées (test d'hypothèse, estimation, ...)

Quand on réalise une analyse de données, il faut bien séparer mentalement ces deux objectifs.

Premiers pas en R

Démarrer rstudio, puis générer un vecteur de 1000 nombres aléatoires distribués selon la loi normale de paramètre (0, 1):

```
rnorm(1000) # premier tirage
rnorm(1000) # nouveau tirage
hist(rnorm(1000)) # afficher un histogramme
```

Interlude : Quelques distributions

```
# Distribution normale
plot(dnorm(seq(-5, 5, .1))) # affiche la densité de proba
pnorm(3) # aire sous la courbe entre -inf et 3
qnorm(.95) # valeur de x telle que P(Z<X)=.95
rnorm(1000) # genere 1000 nombres aléatoires

# Distribution binomiale
barplot(dbinom(0:10, size=10, prob=.5))
barplot(dbinom(0:10, size=10, prob=.2))
pbinom(3, size=10, prob=.5) # proba d'observer 0, 1, 2 ou 3 evènements
qbinom(.95, size=10, prob=.5) # valeur de X telle P(B<=X)=.95
```

Simulation de tirage aléatoire

```
pop = round((scale(rnorm(1000)) * 15 + 180))
hist(pop)
summary(pop)
```

Echantillonner plusieurs fois et calculer la moyenne obtenue, p.ex.:

```
samp1 = sample(pop, 10)
boxplot(samp1)
mean(samp1)
```

```
samp2 = sample(pop, 10)
boxplot(samp1, samp2)
mean(samp2)
```

```
samp3 = sample(pop, 10)
boxplot(samp1, samp2, samp3)
mean(samp3)
```

```
boxplot(samp1, samp2, samp3)
```

```
sampmeans = replicate(100, sample(pop, 10))
hist(sampmeans)
boxplot(sampmeans)
summary(sampmeans)
```

Recommencer avec des tailles d'échantillon plus grande (100).

Recommencer avec une distribution exponentielle:

```
pop = rexp(1000, 2)
hist(pop)
summary(pop)
```

Théorème de la **limite centrale**: Quand n tend vers l'infini, la distribution des moyennes des échantillons de taille n devient de plus en plus distribuée comme une loi normale (centrée sur la moyenne de la population).

Ces exemples fournissent une idée de la variabilité d'échantillons tirés dans une même population. L'écart-type de la distribution de la statistique calculée sur les échantillons est appelé **erreur-standard**. Elle fournit une indication de la précision obtenue avec des échantillons de taille fixée.

Dans la réalité, on ne connaît pas la distribution parente.

Si on en connaît la forme, on peut néanmoins évaluer la précision de la moyenne d'un échantillon de taille n .

Par exemple, pour une proportion, on peut considérer des distributions binomiales (2 événements possibles, situation équivalente au tirage dans une urne, avec remise).

```
sample(0:1, 100, replace=T) # génère 100 tirages pile/face equiprob
sample(0:1, 100, replace=T, prob=c(.3, .7))
rbinom(100, 1, prob=.3)
```

Si on a un échantillon de taille 1000 et que la propriété d'intérêt est présente dans 100 individus, la précision peut être obtenue dans R par:

```
require(binom)
binom.confint(100, 1000, method='exact')

method  x    n mean      lower      upper
1 exact 100 1000 0.1 0.08210533 0.1202879
```

On observe ici que l'intervalle de confiance à 95% est 8.2%–12%. Essayez avec d'autres valeurs.

Attention: ce résultat n'est valide que pour un tirage complètement aléatoire (et une population "infinie"). Dans le cas de plan de sondage plus complexes, il faut utiliser une autre fonction (svyciprop du package survey; cf. exemple pages 53-54 of B. Falissard's *Analysis of Questionnaire Data with R* pour un échantillonnage à deux niveaux).

Intervalles de confiance d'un odd ratio et d'un risque relative

```
require(epi)
examples(twoby2)
```

2 by 2 table analysis:

```
Outcome    : Yes
Comparing  : A vs. B
```

	Yes	No	P(Yes)	95% conf. interval
A	16	10	0.6154	0.4207 0.779
B	15	9	0.6250	0.4218 0.792

		95% conf. interval
Relative Risk:	0.9846	0.6379 1.5197
Sample Odds Ratio:	0.9600	0.3060 3.0117
Conditional MLE Odds Ratio:	0.9608	0.2623 3.4895
Probability difference:	-0.0096	-0.2602 0.2451

Interlude: Bootstrap

Si on ne veut pas faire d'hypothèse sur la forme de la distribution, on peut utiliser une approche de *bootstrap*, qui consiste à se dire que la meilleure estimation de la distribution dans la population est la distribution dans l'échantillon.

Pour obtenir une estimation de la précision de statistiques calculée sur l'échantillon, on effectue des *tirages avec remise* dans celui-ci.

```
data=c(164, 164, 164, 164, 165, 165, 166, 166, 170, 170, 171, 173, 175, 185, 190)
stripchart(data, method='overplot')
# parametric test:
t.test(data)
# bootstrap
n = length(data)
ntirages = 5000 ;
bs = NULL ;

for (i in 1:ntirages) {
  indices = sample(c(1: n), n, replace = T)
  bs[i] = mean(data[indices])
}

hist(bs)
summary(bs)
quantile(bs, c(.025, 0.975))
```

Remarques:

- On peut s'intéresser à beaucoup d'autres statistiques que la moyenne et obtenir des intervalles de confiance (par exemple le minimum, une corrélation, etc...)
- Le package *boot* de R permet de faire cela de manière plus efficace.

Ref: Efron, B. and Tibshirani, R. (1993) *An Introduction to the Bootstrap*. Chapman & Hall.

Statistiques paramétriques et non paramétriques

- Statistiques **paramétriques**: calculs fondés en faisant des hypothèses sur la forme des distributions sous-jacentes.
- Statistiques **non paramétriques**: estimations sans faire d'hypothèses précises sur la forme des distributions.

```
install.packages('car')
```



```

require(car)
data(SLID)
?SLID
str(SLID)
SLID2 = SLID[complete.cases(SLID),]
plot(wages ~ age, data=SLID2)
abline(lm(wages ~ age, data= SLID2), col='blue', lwd=2)
lines(lowess(SLID2$age, SLID2$wages, f=0.1), col='red', lwd=2)
smoothScatter(SLID2$age, SLID2$wages)
bin = hexbin(SLID2$age, SLID2$wages, xbins=50)
par(plot='NEW')
plot(bin, main="Hexagonal Binning")
abline(lm(wages ~ age, data= SLID))

```

Le **Big Data** est un champ d'application des méthodes non-paramétriques.

Comparaison de deux populations

Supposons qu'on tire des échantillons dans deux populations pour comparer celles-ci, par exemple avec un test de T (test de Student). On peut refaire des stimulations.

```

pop1 <- rnorm(1000000, mean=180, sd=15)
pop2 <- rnorm(1000000, mean=185, sd=15)
t.test(sample(pop1, 1000), sample(pop2,1000), var.equal=T)

```

On pourrait, par essais et erreurs, déterminer la taille de l'échantillon nécessaire pour détecter une différence significative (au seuil de 5%) dans 80% des cas.

Mais une fonction R permet cela automatiquement:

```
power.t.test(delta=5, sd=5, sig.level=.05, power=.8, type='two.sample')
```

Précision et représentativité d'un échantillon

Idéalement, pour permettre de généraliser, l'échantillon doit être :

- **précis** : d'une taille suffisante pour que l'erreur d'estimation qu'il introduit soit acceptable.
- **représentatif** : sa composition doit être semblable à celle de la population-mère.

Echantillonnage complètement aléatoire

- Il s'agit d'un tirage parfaitement aléatoire dans toute la population.

Celle-ci ayant une taille N , chaque individu à une probabilité $1/N$ d'être inclu.

- Il faut disposer de la liste complète des membres de la population-mère pour pouvoir mettre en oeuvre une véritable sélection aléatoire (et d'une fonction d'extraction aléatoire comme `sample`).

Si l'échantillon ne respecte pas ce critère de représentativité, il est considéré comme *biaisé* et il faut envisager d'effectuer un **redressement** (du moins *si l'on veut absolument des estimateurs non biaisés*).

*Redressement d'échantillon**Le redressement par suppression*

Afin de retrouver les proportions attendues (celles de la population-mère), on supprime aléatoirement des répondants parmi les catégories sur-représentées.

Cela entraîne la réduction de la taille de notre échantillon, ce qui est frustrant vu les efforts réalisés pour motiver les personnes contactées à répondre et vus les coûts engendrés. Par ailleurs, on va perdre en précision puisque l'erreur associée va augmenter.

Cette stratégie peut être néanmoins une bonne solution dans un protocole de collecte par Internet qui permet de contacter rapidement et à moindre coût un grand nombre d'interlocuteurs. On peut ainsi extraire après coup et selon une méthode aléatoire, un échantillon représentatif selon des quotas pré-définis.

Application: Redressement par suppression dans R

```
require(car)
data(SLID)
str(slid)
table(SLID$sex)
n = min(table(SLID$sex))
males = subset(SLID, sex=='Male')
nrow(males)
females = subset(SLID, sex=='Female')
nrow(females)
females = females[sample(1:nrow(females), n), ]
nrow(females)
```

Le redressement par pondération

On conserve toutes les réponses enregistrées mais on attribue à chaque répondant un « poids » particulier en fonction de la catégorie à laquelle il appartient.

Par exemple, si il y deux fois moins de femmes que prévu dans l'échantillon, le « poids » d'une femme sera 2 et la réponse de chaque femme comptera double.

Voir <http://www.suristat.org/article27.html>

Pondérer ou non: un choix parfois délicat

Considérons une entreprise avec la structure de salaire suivante:

Sexe	Status	Salaire moyen	Effectif dans l'entreprise
Femme	cadre	3000	2
Femme	non-cadre	2000	50
Homme	cadre	3000	8
Homme	non-cadre	2000	50

Question: les hommes et les femmes perçoivent-ils le même salaire?

Application: Pondération dans R

```
sal = c(3000, 2000, 3000, 2000)
sex = c('F', 'F', 'H', 'H')
status = c('c', 'n', 'c', 'n')
eff = c(2, 50, 8, 50)
data.frame(sex, status, sal, eff)
# salaire moyen dans l'entreprise
mean(sal)
weighted.mean(sal, eff)
# salaire par sexe
# non pondéré
for (s in unique(sex)) { print(paste(s, mean(sal[sex==s]))) }
tapply(sal, sex, mean)
# pondéré
for (s in unique(sex)) {
  print(paste(s, weighted.mean(sal[sex==s], eff[sex==s]))) }
}
```

Note: Le package survey fournit des fonctions statistiques qui prennent systématiquement en compte des poids.

Echantillonnage stratifié (méthode des quotas)

La stratégie de redressement peut être évitée si on prend en compte la structure de la population dès le plan de sondage.

La population est découpée en groupes (par exemple, groupes d'âge, de sexe, de niveau socio économique) dont on connaît les effectifs.

On prélève alors des échantillons aléatoires à l'intérieur de chaque groupe.

Si les tailles des échantillons respectent les proportions des groupes dans la population: on a une meilleure précision que dans le cas du tirage complètement aléatoire.

Par exemple: pour comparer les tailles des français et des allemands, on peut constituer des échantillons qui respectent les proportions d'hommes et de femmes dans chaque catégorie d'âge.

Pour en savoir plus, voir Thomas Lumley *Complex Surveys: A guide to analysis using R* (package survey)

Echantillonnage par "grappes" (cluster sampling)

On effectue d'abord un tirage aléatoire d'unités plus grandes (p.ex. villes, ou des écoles). Puis on mesure tout ou parti des individus dans ces unités.

- Avantage: Il est plus facile et moins coûteux à mettre en oeuvre que les méthodes précédentes
- Désavantages:
 - Plus d'individus sont nécessaires pour obtenir une précision identique à l'échantillonnage aléatoire complet.
 - Il faut tenir compte de la dépendance entre les individus d'un même "cluster". Cela complique nettement les analyses statistiques (on doit utiliser des modèles-hiérarchiques)

Déterminer la taille des échantillons pour estimer une proportion

Dans le cas d'un tirage purement aléatoire, on doit fournir une estimation de la fréquence attendue (p), et la demi-largeur de l'intervalle de confiance à 95% (δ):

```
require(epiDisplay)
n.for.survey(p=0.1, delta=0.02)
```

Une formule approximative permet d'évaluer la taille de l'échantillon:

$$n = \frac{p(1-p)}{(e/2)^2}$$

(remarque: le maximum de $p(1 - p)$ est de 0.25)

A retenir: **la précision est proportionnelle à la racine carrée de la taille de l'échantillon**

Pour les échantillonnages par strate ou par grappe, les formules sont nettement plus compliquées.

Pour un échantillonnage par grappe, il faut tenir compte du fait que les données des deux individus dans le même groupe sont corrélées. On introduit la notion d'effet de design (Falissard, p.72):

```
n.for.survey(p=0.1, delta=0.2, deff= 4)
```

Déterminer la taille d'un échantillon pour comparer deux groupes

```
require(epiDisplay)
```

```
# Comparer 2 proportions:
```

```
n.for.2p (p1, p2, alpha = 0.05, power = 0.8, ratio = 1)
```

```
# Comparer 2 moyennes:
```

```
n.for.2means (mu1, mu2, sd1, sd2, ratio = 1, alpha = 0.05, power = 0.8)
```

Arguments:

p: estimated probability

delta: difference between the estimated prevalence and one side of the 95 percent confidence limit (precision)

popsiz: size of the finite population

deff: design effect for cluster sampling

alpha: significance level

mu1, mu2: estimated means of the two populations

sd1, sd2: estimated standard deviations of the two populations

ratio: n2/n1

Remarque sur les objectifs des analyses statistiques

1. **Test d'hypothèse** : une variable influence-t-elle une autre?
2. **Estimation** : quel est la taille de l'effet de X sur Y (dans les conditions Z, W...) ?
3. **Prédiction** : à partir des caractéristiques de nouveaux individus, peut-on prédire la valeur des variables dépendantes (p.ex. probabilité d'être fumeur) ?

Ex: si on veut tester si les droitiers réagissent plus rapidement avec la main droite qu'avec la main gauche, on peut se contenter de faire l'expérience sur des étudiants d'université: si l'effet existe, ils est raisonnable qu'il soit présent chez tous les humains. Mais l'amplitude de l'effet peut dépendre de l'âge.

Dans des approches exploratoires, un test biaisé n'est pas un problème car l'important est de détecter un effet.

Si l'on désire tester la valeur prédictive des modèles, on peut utiliser des approches de **cross-validation**.

Partie 3: Construction d'un questionnaire

Les qualités fondamentales d'un questionnaire:

- l'outil choisi mesure bien ce pour quoi il a été construit (**validité**)
- les mesures réalisées sont bien reproductibles (**fiabilité** (reliability)) et permettent de détecter les effets recherchés (puissance statistique)

Vérification de la validité

Lors de la phase de définition des hypothèses générales de l'enquête ainsi que de ses objectifs, il est nécessaire de bien préciser l'information désirée.

"construct validity": Est-ce que les définitions opérationnelles choisies capturent bien les concepts (p.ex. anxiété) pertinents ? Ne sont-elles pas influencées par des facteurs non pertinents ?

Exemple: si je veux mesurer l'intelligence par le QI, et que le test de QI que je construis est très long, je peux être en train de mesurer la persévérance plutôt que l'intelligence. (Au passage, certains questionnent si la mesure de QI est une bonne mesure de l'intelligence).

Lors de la conception du test, il faut jouer l'avocat du diable et chercher des explications alternatives. Tant qu'il y en a, modifier le test.

Introduire plusieurs instruments de mesures (pour une enquête, plusieurs questions), dont certaines utilisant des méthodes "éprouvées" (si elles existent), et vérifier a posteriori que leurs résultats corrélaient bien.

Comment vérifier la fiabilité?

cross-validation (split-half, k-fold, leave-one out)

Retest sur un nouveau groupe indépendant.

Inter-rater reliability (IRR)

S'il y a des réponses qui doivent être classifiées par un humain, il faut vérifier si un autre humain donne des résultats similaires.

Exemple: deux juges classent les réponses des participants dans deux catégories:

Juge1 Rep1 Rep2

Juge2		
Rep1	a	b
Rep2	c	d

Pour évaluer le degré d'accord, on calcule le Kappa de Cohen.

Par exemple, si deux psychiatres évaluent 23 patients pour décider s'ils ont un tristesse douloureuse. Les données (accordqual.csv) sont disponibles au format csv2 sur la page <http://hebergement.u-psud.fr/biostatistiques/#livre&id=02&r=partie02>

```
require(psy)
data = read.csv2(' http://hebergement.u-psud.fr/biostatistiques/includes/livre/02/partie02/fichiers/accor
table(data)
ckappa(data)
```

Puissance statistique (sensibilité):

- *au moment de la conception:* s'assurer qu'on a un nombre adéquat de mesures (pour cela, il faut estimer les tailles d'effets et leurs variances)
- *après-coup:* à partir des effets effectivement observés et leur variabilité, on peut évaluer le nombre de sujets qu'il aurait fallu pour détecter une différence avec une certaine probabilité.

Critères d'inclusion/exclusion

Il peut être nécessaire de définir des critères d'exclusion.

Par exemple, en imagerie cérébrale sur le sujet sain, j'évite les personnes qui consomment trop d'alcool ou de stupefiants, les femmes potentiellement enceintes, etc.

Cela fait parti de la définition de la population-mère.

Différents modes de recueil des données

- Questionnaires auto-administrés, par courrier ou Par Internet (site ou email)
- Interview téléphonique

Avantage: si les interviewers sont bons, la qualité des réponses peut être meilleure que dans les questionnaire auto administrés.

Désavantages: durée limitée ne permettant que des questions assez simples.

- Interview en face-à-face

- interview structurée: les mêmes questions sont posées exactement de la même manière à tous les participants
- interview non structurée: interaction libre.
- Groupe de discussion

Petit groupe de personnes partageant des caractéristiques (e.g. femmes entre 30 et 40 ans, ...). L'interviewer est un "facilitateur" qui pose des questions ouvertes.

Intéressant pour prétester les questions d'une enquête afin de choisir les plus pertinentes.

Construction pratique d'un questionnaire

Des logiciels permettent de construire des interfaces avec des *masques de saisie* pour les différentes questions, permettant de spécifier des conditions de validité des réponses, d'effectuer des branchements en fonction des réponses, ...

Des modules permettent : - la saisie des données - leur analyses statistique ou leur exportation dans des fichiers (tables, bases de données, ...) exploitable par des logiciels comme R, SAS, SPSS, Stata...

Offline

- epidata (<http://www.epidata.dk/>) puissant, multiplateforme (Windows, Mac, Linux), gratuit. (cf. TP sur <http://hebergement.u-psud.fr/biostatistiques/#livre&id=01&r=partie01>)

Online

- google forms <https://docs.google.com/forms> très pratique pour des questionnaires simples (p.ex. ne permet pas de branchement). Fourni les résultats sous forme de feuille de calcul google
- surveymonkey <https://fr.surveymonkey.com/>: simple à utiliser, puissant, 400€/an (900€ pour la version plus puissante).
- limesurvey: <https://www.limesurvey.org/> libre, très puissant, peut être installé sur son propre serveur web.
- Moodle (plateforme de cours en ligne) utilise un format GIFT pour construire les questionnaires.

Les questions

- Quand elles portent sur des individus, elles permettent d'obtenir trois types de données:
- **Factuelle.** Ex: caractéristiques démographiques, fumeur non fumeur, ...

- **Comportementales:** Ex: style de vie (Ex: prenez vous souvent le metro?)
- **Attitudinale:** opinions, valeurs, croyances (pensez-vous que 'X' soit vrai/faux)
- Poser de bonnes questions est un art.

Voir A. N. Oppenheim (1992) *Questionnaire design, Interviewing and Attitude Measurement Continuum*.

"The world is full of well-meaning people who believe that anyone who can write plain English and has a modicum of common sense can produce a good questionnaire. This book is not for them."

- Questions à choix fermé vs. questions à choix ouvert.

Exemples de questions à choix fermé

Likert scale

Strongly disagree | Disagree | Neither agree nor disagree | Agree | Strongly Agree

Echelle sémantique différentielle:

This test is:

difficult ____:____:____:____:____:____:____ easy
useless ____:____:____:____:____:____:____ useful

Choix multiple:

During French class, I would like:

- (a) to have a combination of French and English spoken
- (b) to have as much English as possible spoken
- (c) to have only French spoken

Questions à choix ouvert

A utiliser de préférence pour demander au sujet d'apporter des clarifications

Exemples de questions à choix ouvert:

Quel est votre programme télé favori?

Si vous avez mis une note inférieure ou égale à 2 à votre séjour, merci d'expliquer brièvement vos raisons

Questions guidées:

Une chose que j'ai apprécié est:

Une chose que je n'ai pas apprécié est:

Types de Variable

- **Nominale** : la variable peut prendre des valeurs dans un ensemble fini de catégories distinctes. Ex: Vrai/Faux. Couleur de cheveux.
- **Ordinale** : les catégories peuvent être ordonnées. Ex:
- **intervalle** : la différence numérique entre les catégories a du sens. Ex: température (la différence entre 40 degré C et 20 degré C est la même qu'entre 20 et 0, mais 40 n'est pas deux fois plus chaud que 20 degrés)
- **ratio**: le ratio entre deux valeurs a du sens. Ex: poids (4 kg est 2 fois plus lourd que 2 kg)

Les procédures statistiques dépendent du type de données. Par exemple :

- Pour comparer deux proportions (variable nominale à deux valeurs), on pourra utiliser un test de chi2 (e.g. `chisq.test`), une régression logistique ou un modèle loglinéaire.
- Pour comparer deux moyennes de variables de type intervalle ou ratio, on utilisera un test de Student ou de Welsh.

Interlude: rapports de cote et risque relatif

Une probabilité p est un nombre entre 0 et 1.

Pour comparer deux probabilités, plutôt que de regarder leur différence, on utilise plutôt le *risque relatif* ou le *rapport de cote*.

Le risque relatif est simplement le ratio:

Par exemple, Si 10 % des fumeurs ont eu un cancer du poumon, et 5 % des non-fumeurs ont eu un cancer du poumon, le risque relatif (RR) est $.10/.05 = 2$.

On peut aussi transformer les probabilités en cotes ("odds" en anglais), par exemple "10 contre 1", et on utilise des rapports de cotes (odd ratio) pour comparer des probabilités.

$$c = \frac{p}{1-p}$$

```
odd = function(p) { p / (1 - p) }
plot(odd(seq(0, 1, by=.01)), type='l')
odd(.1)/odd(.05)
```

2.111111

On peut obtenir des intervalles de confiance pour le RR et les odd ratio avec la fonction `twoby2` du package `Epi`

```
require(Epi)
examples(twoby2)
```

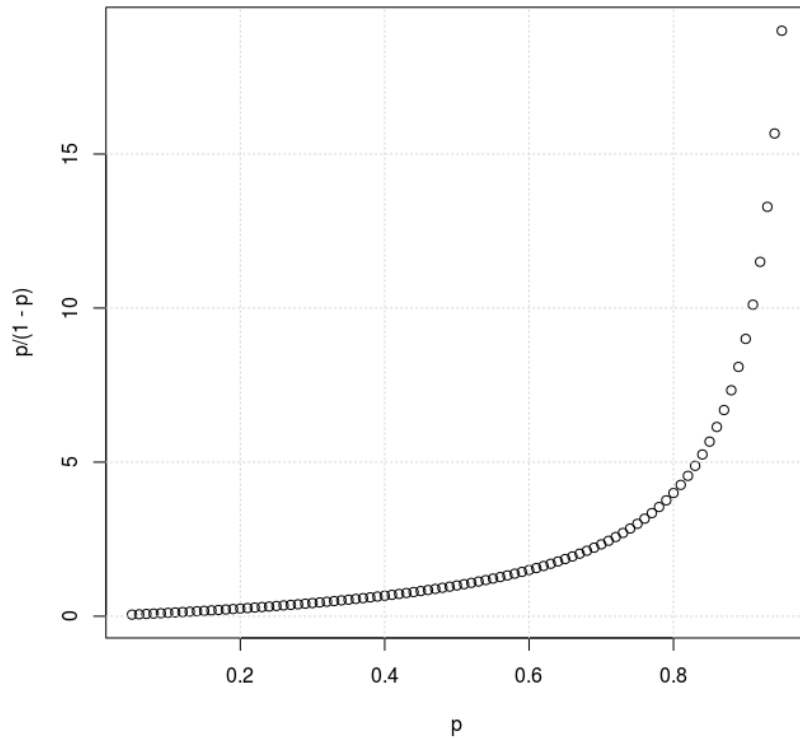


Figure 12: Conversion de probabilité en côte

Conseils pratiques pour des questionnaires auto-administrés

Afin de maximiser le taux de réponse:

- Les questions doivent être *simples*.
- les questions doivent suivre un ordre logique,
- La concision doit l'emporter sur la tentation d'être exhaustif.
- Faire une mise en page soignée. Aérer le texte. Format petit livret préférable à pages A4.
- Eviter d'avoir un saut de page au milieu d'un item.
- Idéalement, pas plus de 4 pages, 30 minutes max.
- mettre les questions factuelles à la fin du questionnaire

Contrôle de qualité

Nettoyage des données

- Données absurdes Ex: Si code numérique devait être entre 1 et 7, une valeur en dehors de cet intervalle est "absurde"; il faut la remplacer par un 'NA' (donnée manquante).
- Réponses contradictoires à des items distincts (il faut envisager d'éliminer ces questionnaires)

- Données implausibles (outliers). C'est le cas problématique. Cela reflète-t-il la vérité ou une erreur? Approches possibles:
- utiliser des statistiques robustes (p.ex. médiane plutôt que la moyenne).
- faire l'analyse avec et sans les outliers, et si le résultat ne dépend pas de ceux-ci, tout va bien.
- Données manquantes:
- supprimer l'enregistrement complet. Ou ne pas l'utiliser dans les analyses qui ont besoin de cette variable.
- interpoler avec les données des participants similaires (imputation)

Construction et validation psychométrique d'un questionnaire.

Pour un bon questionnaire, on s'attend essentiellement à ce que :

- les items doivent avoir suffisamment de variance (si toutes les réponses à un item sont bloquées sur la même valeur, cet item n'apporte pas d'information)
- les items censés mesurer le même "construct" doivent corrélérer plus entre eux qu'avec les autres items.

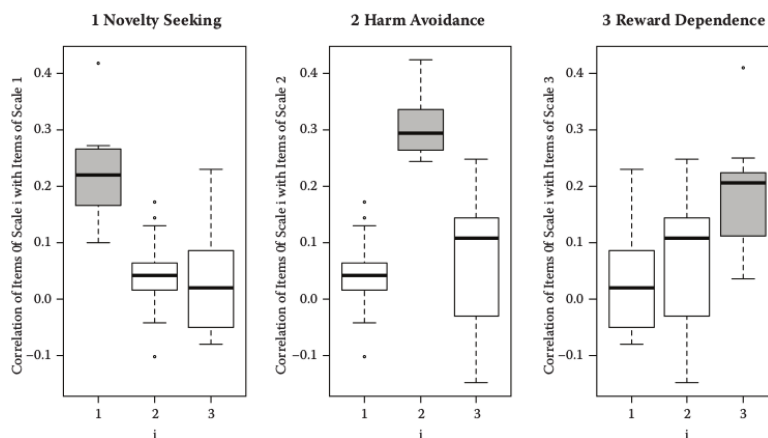


FIGURE 7.5

Distribution of inter-item correlations for the temperament subscales. The different items correlate better with other items from the same subscale (grey boxes are above white boxes). These results can be related to those presented in Figure 7.3.

Remarques:

- On utilise également l'Analyse Factorielle (Falissard, Chap.7)
- Les graphiques de Bland-Altman sont aussi pertinents.
- Prévoir (au moins) une phase de pilotage où une version préliminaire du questionnaire est distribuée puis analysée de manière à ne conserver que les items les plus efficaces.

Traçabilité des données

Identificateur unique sur chaque questionnaire

- Le générer lors de la production des questionnaires
- Si le questionnaire est administré par ordinateur, une bonne idée consiste à générer un identifiant sur la base de la date/heure/identifiant de l'ordinateur (IP).

Une fois les données recueillies et saisies dans un logiciel :

- conserver la forme papier éventuelle dans des archives pour vérifier d'éventuelles erreurs de saisies (données aberrantes).
- conserver la forme informatique sous forme de fichier TEXTE (Interdiction absolue de conserver les données dans un format binaire propriétaire (feuille EXCEL, base de données ACCESS). Privilégier des formats lisibles par l'humain (json, csv, plutôt que XML).

Sécurité:

- calculer une somme de contrôle md5 pour chaque fichier de données et les conserver dans un fichier qui ne peut pas être modifier. Cela permettra de vérifier plus tard l'intégrité des données.
- protéger les données brutes: au minimum, zip protégé par mot de passe lors d'envois par email ou de partage sur des serveurs Internet (Dropbox, Google Drive, ...) — sauf si les données brutes de l'enquête sont destinées à être rendues publiques.

Anonymisation

Si les données d'une enquête doivent être partagées, il est presque toujours nécessaire d'**anonymiser** celles-ci, c'est à dire de supprimer des champs tels que nom, adresse, et toute information qui pourrait permettre d'identifier les participants.

Le mieux est de faire cela au moment du codage des données.

Si on n’a jamais besoin de revenir à l’individu, jeter purement et simplement ces données.

Si on risque d’avoir besoin de remonter à l’individu — p.ex. enquête longitudinale — garder dans un fichier séparé un table reliant les identifiants anonymes et les informations identifiantes.

Traçabilité des traitements

AUCUNE intervention manuelle n’est autorisée sur les données !

TOUS les traitements (filtrage, ré-organisation des données, calculs statistiques) doivent être *automatisés*, par exemple avec des scripts écrit en langage R.

Les outils de “literate programing” comme *rmarkdown* (logiciel R) ou les *notebooks jupyter* (Python) sont très utiles pour documenter une analyse de données et produire des rapports. Notamment cela évite les copier-coller entre les sorties du logiciels d’analyse et les rapports.

A faire: Démonstration de Rmarkdown.

Partie 4: Manipulation et Exploration des données

Importation des données

Les données, sous forme de tables, sont lues dans des `data.frame`, typiquement avec les fonctions

```
read.csv(filename)
read.table(filename)
```

La librairie `foreign` permet de lire les formats Minitab, S, SAS, SPSS, Stata, Systat, Weka, dBase,

On peut se faire une idée du contenu d’un `data.frame` `dataf`.

```
str(dataf)
head(dataf)
names(dataf)
```

On peut accéder au contenu d’une colonne d’un `data.frame` avec la syntaxe `dataf$colname`.

Manipulations

```
subset
merge
```

Recodage

```
is.na
```

```
cut
ifelse(test, value-if-yes, value-if-no)
require(car)
?recode
```

Examiner des distributions

- Variables discrètes

```
table(x)
barplot(table(x))
```

- Variables continues

```
summary(x)
stem(x)

stripchart(x, method='stack')
boxplot(x)

plot(density(x))
rug(x)
```

Examiner des relations

- variables discrètes:

```
table(x, y, z)
ftable(x, y, z)
xtabs(~ x + y + z)
mosaic.plot(table(x, y))
chisq.test()
```

- variables continues:

```
plot(x, y)
require(car)
scatterplot(x, y)
smoothScatter(SLID2$age, SLID2$wages)

cor.test(x, y)
t.test(x, y)
```

More than 2 variables

```
plot(x, y, col=z)
pairs(cbind(x, y, z))
scatterplot.matrix(cbind(x, y, z))
coplot(x ~ y | a + )

lm(z ~ x + y) # regression multiple
```

R graphics Gallery

<http://www.r-graph-gallery.com/>

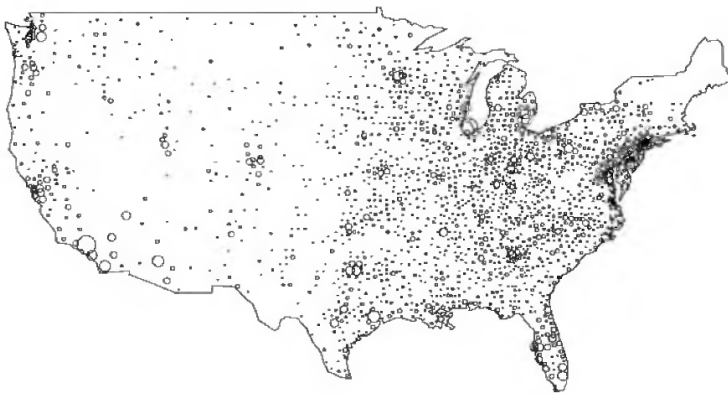
Cartes géographiques

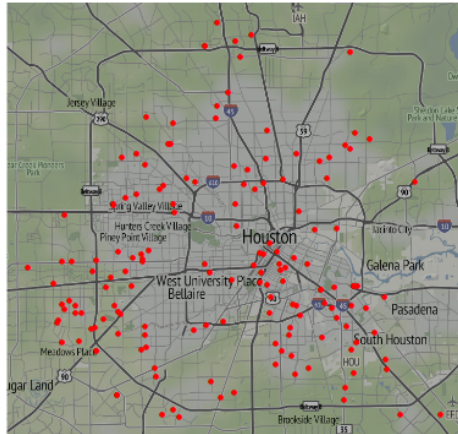
Figure 3.1 A simple random sample of 10,000 voter locations from the USA. Circles are at the center of each county, with area proportional to the number sampled. The largest sample is 257, in Los Angeles County. County-level data for some areas of New England were not available.

- *ggmap: Spatial Visualization with ggplot2* par David Kahle and Hadley Wickham <https://journal.r-project.org/archive/2013-1/kahle-wickham.pdf>
- *Géocoder en masse avec R et sans Google Maps* par Timothée Giraud <http://rgeomatic.hypotheses.org/622>
- *Plotly Scatter Plots on Maps in R* <https://plot.ly/r/scatter-plots-on-maps/>

Données à analyser

- *caith* Colours of Eyes and Hair of People in Caithness (package MASS)


```
murder <- subset(crime, offense == "murder")
qplot(lon, lat, data = murder, colour = I('red'), size = I(3), darken = .3)
```



- A 4 by 5 table with rows the eye colours (blue, light, medium, dark) and columns the hair colours (fair, red, medium, dark, black).
- SLID Survey of Labour and Income Dynamics (package car)
 - wages Composite hourly wage rate from all jobs.
 - education Number of years of schooling.
 - age in years.
 - sex A factor with levels: 'Female', 'Male'.
 - language A factor with levels: 'English', 'French', 'Other'.
- mtcars Motor Trend Car Road Tests

The data was extracted from the 1974 *Motor Trend* US magazine, and comprises fuel consumption and 10 aspects of automobile design and performance for 32 automobiles (1973-74 models).
- données du livre de Bruno Falissard *Comprendre et utiliser les statistiques dans les sciences de la vie*: <http://hebergement.u-psud.fr/biostatistiques/#livre&id=02&r=partie02>

Références

- Général:
 - Lohr, S. (2010) *Sampling Design and Analysis*. Brooks/Cole.
 - Oppenheim, A. N. (1992) *Questionnaire Design, Interviewing and Attitude Measurement*. Continuum.
- Analyses avec R:

- Zuur, Ieno Meester (2009) *A Beginner's Guide to R* Springer
- Falissard, B. (2012) *Analysis of Questionnaire Data with R*. CRC Press.
- Lumley, T. (2010) *Complex Surveys: A guide to analysis using R*. Wiley (package survey)
- Chongsuvivatwong V. (2013) *Analysis of epidemiological data using R and Epicalc* https://cran.r-project.org/doc/contrib/Epicalc_Book.pdf (note: le package epicalc a été renommé epiDisplay)
- Aides mémoire: <https://www.rstudio.com/resources/cheatsheets/>