

Plan de sondage informatif : comment le définir et en tenir compte ?

Séminaire en ligne sur les sondages du groupe
Enquêtes de la Société Française de Statistique
organisé en collaboration avec l'Université de
Neuchâtel

20 novembre 2025

D.Bonnéry

1

La sélection est informative
lorsqu'il convient d'en tenir
compte

1. La sélection est informative lorsqu'il convient d'en tenir compte



Pfeffermann, Danny, Krieger, AM Abba M AM, & Rinott, Yosef. 1998.

Parametric distributions of complex survey data under informative probability sampling.

Statistica Sinica, 8(4), 1087–1114.

Les données d'enquête sont le produit de deux processus aléatoires.

- Le processus dont est issu la population d'intérêt et ses caractéristiques
- Le processus de sélection

La sélection est informative quand il convient d'en tenir compte lors de l'inférence sur le modèle de population à partir des données d'enquête.

1. La sélection est informative lorsqu'il convient d'en tenir compte

Lorsque les observations sont le produit

- d'un processus d'intérêt
- d'un processus de nuisance.

le processus de nuisance est informatif lorsqu'il convient d'en tenir compte lors de l'inférence.

Processus de nuisance:

- erreur de mesure
- caviardage
- censure
- non réponse
- nombre de tirages par unité
- taille de la population
- taille de l'échantillon
- étiquetage

1. La sélection est informative lorsqu'il convient d'en tenir compte

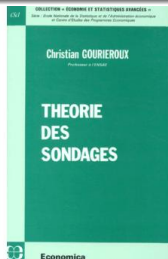
- Modèle de la variable d'intérêt: $V \sim P$;
- Tout le reste est nuisance: $\bar{V} \mid V \sim P'$

Observation = $X = \text{function}(V, \bar{V})$.

2

Modèle générique pour les
sondages et sélection
informatrice dans la littérature
scientifique

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

II C - LE MODELE DE LA THEORIE DES SONDAGESII C 1 - MISE EN EVIDENCE D'UN MODELE
STATISTIQUE PARAMETRIQUE

Dans ce paragraphe, nous supposons donné le plan de sondage Π . La loi des observations (s, Y_s) est l'image de Π par l'application : $s \longrightarrow (s, Y_s)$. Nous notons $\Pi^I Y$ cette loi. Elle dépend de la valeur inconnue du caractère Y . Le modèle expliquant les observations peut ainsi être considéré comme un modèle statistique paramétrique, de paramètre Y . Nous pouvons donc mener l'inférence sur ce paramètre en utilisant les méthodes habituelles de Statistique Mathématique.



Gourieroux, C. 1981.

Théorie des sondages.

La distribution des observations sous le modèle fixe de population pour l'inférence basée sur le plan appartient à un **modèle paramétrique**, les caractéristiques de la population étant le paramètre.

Pour une population de taille N , un plan de sondage P^S , le modèle

- de population est: $P^Y \in \{\text{dirac}_\theta \mid \theta \in \Theta = \mathbb{R}^N\}$
- des observations est: $P^{S,Y[S]} \in \{(P^S)^{x \rightarrow \theta[x]} \mid \theta \in \Theta = \mathbb{R}^N\}$

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Godambe, V. P. 1966.

A New Approach to Sampling from Finite Populations. I
Sufficiency and Linear Estimation.

Journal of the Royal Statistical Society, 28(2), 310–319.



Ericson, W . A . 1969.

Subjective Bayesian Models in Sampling Finite Populations.

Journal of the Royal Statistical Society, 31(2), 195–233.

Les auteurs expliquent que dans le cas du modèle de population fixe, la vraisemblance est "uninformative".

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Scott, Alastair J. 1975.

Some comments on the problem of randomisation in surveys.

In: Int. Assoc. of Survey Statisticians, 2nd Meeting, Warsaw, September 1975.



On the problem of randomization in survey sampling.

In: Sankhya, 39, 1977.

“

Scott (1977)

”

[If the] posterior distribution depends only on the sample actually drawn and not on the sampling design used to draw it, provided that [the sample of the population elements] and [the study variable on the population] are independent random variables. (If this condition of independence is satisfied the design is said to be non-informative.)

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

- ❶ Premier article (parmi ceux que j'ai lu) qui proposent une définition d'un plan de sondage non informatif.
- ❷ Autre aspect important: la définition n'est pas figée à un cadre particulier, mais s'applique à des cas différents, ou les observations peuvent contenir plus ou moins d'information (variables auxiliaires sur l'échantillon, la population...).

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Rubin, DB Donald B. 1976.

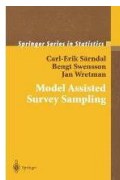
Inference and missing Data.

Biometrika, **63**(3), 591–592.

- ❶ Ne définit pas ce qu'est un processus informatif en général, mais, dans un cas particulier
- ❷ Définit ce qu'ignorer un processus de nuisance particulier (la non réponse) signifie
- ❸ Énonce des conditions pour que la vraisemblance lorsque le processus de nuisance est ignoré est la même que la vraisemblance lorsqu'il ne l'est pas.
- ❹ La vraisemblance est la vraisemblance conditionnelle au processus de nuisance.

Le terme ignorable n'apparaît pas, mais le terme "at random" oui.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Cassel, CM, Sarndal, CE, & Wretman, JH.
1977.

*Foundations of inference in survey
sampling.*

Définition tautologique limitée au modèle de population fixe: le plan est non informatif lorsque les probabilités d'inclusion et les variables d'intérêt sont indépendantes, ce qui est toujours le cas dans un modèle de population fixe.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Little, Roderick J. A. RJA. 1982.

Models for Nonresponse in Sample Surveys.

Journal of the American Statistical Association, **77**(378),
237–250.

Dans la lignée de Rubin, 1976, donne la définition d'un plan de sondage ignorable.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



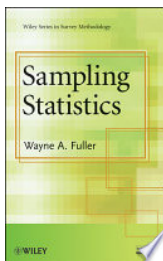
Sugden, RA, & Smith, T. M. F. 1984.

Ignorable and informative designs in survey sampling inference.

Biometrika, 71(3), 581–592.

- Premier article (à ma connaissance) à opposer to plan ignorables et informatifs dans un article.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Fuller, Wayne A. 2011.
Sampling statistics.
Vol. 560.
Hoboken, NJ: Wiley.

“

If $E[x_i \pi_i e_i] \neq 0$, it is sometimes said that the design is “informative” for the model.

”

Fuller (2011)[p. 355 sec. 6.3.2]

Les objets x_i , π_i et e_i sont la variable d'intérêt, la probabilité d'inclusion, le résidu de la régression linéaire de x_i sur la variable de sondage.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

Une définition contradictoire

Le sondage est informatif $\Leftrightarrow$ dépendance entre π_k et Y_k après conditionnement sur les covariables du modèle X_k .

Ce n'est

- ni suffisant,
- ni nécessaire.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

Plan informatif: $\pi \perp Y$

Sondage par grappe:

- $(Y_k)_{k \in U} \sim \text{iid } \mathcal{B}(1/2)$,
- $J = Y + \varepsilon$,
- Grappe U_0, U_1 obtenues après tri selon J ,
- Probabilité de sélection d'un cluster : $\frac{1}{2}$,
- $\pi_k = \frac{1}{2}$,
- $(Y_k)_{k \in S}$ ne sont pas iid $\mathcal{B}(1/2)$.

Plan non-informatif: π et Y sont dépendants

- $(Y_k)_{k \in U} \sim \text{iid } \mathcal{B}(1/2)$,
- $\pi_k = \begin{cases} 1/2 & \text{if } Y_k = 0 \\ \begin{cases} 1/4 \text{ w.p. } 1/2 \\ 3/4 \text{ w.p. } 1/2 \end{cases} & \text{if } Y_k = 1, \end{cases}$
- Plan: $\text{Poisson}(\pi)$,
- $(Y_k)_{k \in S}$ sont iid $\mathcal{B}(1/2)$.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Pfeffermann, Danny, Krieger, AM Abba M AM, & Rinott, Yosef. 1998.

Parametric distributions of complex survey data under informative probability sampling.

Statistica Sinica, 8(4), 1087–1114.

D. Pfeffermann introduit les notions de distribution de population et d'échantillon, ces distributions pouvant différer en présence de sélection informative.

- La distribution d'échantillon est une loi dominée par la distribution de population
- La dérivée ρ de Radon Nikodyn est la probabilité d'être sélectionnée ou l'espérance des probabilités de sondages conditionnellement à la caractéristique d'intérêt.
- de la connaissance ou modélisation de ρ peut être dérivée une vraisemblance approchée.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

$$\rho(y) = [E[I_k \mid Y_k = y]],$$

Approximativement : $(Y_k)_{k \in S} \stackrel{\text{i.i.d.}}{\sim} \rho_Y$

- Intérêt pratique pour des modèles de superpopulation avec résidus par unités indépendants.
- Estimateurs plus efficaces que les estimateurs de pseudo vraisemblance.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique



Daniel Bonnéry, Breidt, F. J., and Coquet, F. (2012).
Uniform convergence of the empirical cumulative distribution
function under informative selection from a finite population.
Bernoulli Journal, 2012.



Daniel Bonnéry, F. Jay Breidt, and François Coquet.
*Kernel Estimation for a Superpopulation Probability Density
Function under Informative Selection*
Metron, 75(3), 2017.

<http://rdcu.be/wtPL>.



Daniel Bonnéry, F. Jay Breidt, and François Coquet.
Asymptotics for the maximum sample likelihood estimator
under informative selection from a finite population.
Bernoulli, 11(2):655–679.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.1. Notation conventions and definitions

2.1. Notation conventions and definitions

- Généralisation des définitions de processus ignorables.



Daniel Bonnery and Joseph Sedransk, 2020.

On the definition of informative vs. ignorable nuisance process.

arXiv, **math.ST**1906.02733.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.1. Notation conventions and definitions

“

it is inevitable, [...] that a paper dealing with [the] problem [of what conclusions regarding sufficient statistics may be drawn from the existence of uniformly most powerful tests, or vice versa] should bear some mark of the theory of functions, in spite of its concern with statistical questions

”

Neyman & Pearson (1936)

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.1. Notation conventions and definitions

2.1.1. Théorie des ensembles

- Soit $h : \mathcal{E} \rightarrow$, $\sim_h$ est la relation d'équivalence sur $\mathcal{E}$:
 $x \sim_h x' \leftrightarrow h(x) = h(x')$. Les classes d'équivalence sur $\mathcal{E}$ pour $\sim_h$ sont notées $\mathcal{E}/h$, la classe de x pour $\sim_h$ est $\text{class}_h(x) = h^{-1}(\{h(x)\})$.
- $h : \mathcal{E} \rightarrow$ dépend déterministiquement $h' : \mathcal{E} \rightarrow$ seulement (on note $h < h'$) si et ssi $\exists g : \rightarrow$ tel que $h = g \circ h'$.
- le complément de $h : \mathcal{E} \rightarrow$ est toute fonction $\bar{h} : \mathcal{E} \rightarrow$ telle que $(\text{class}_h, \text{class}_{\bar{h}}) : \mathcal{E} \rightarrow (\mathcal{E}/h) \times (\mathcal{E}/\bar{h}), x \mapsto (\text{class}_h(x), \text{class}_{\bar{h}}(x))$ est injective. Si une fonction $\bar{h}$ satisfait cette propriété, on écrit $(h, \bar{h}) \sqsupseteq \mathcal{E}$. Pour un complément $\bar{h}$ de h , on définit la fonction $\sqcap_{h, \bar{h}} : \{h(x), \bar{h}(x) : x \in \mathcal{E}\} \rightarrow \mathcal{E}$ comme l'inverse de $(\text{class}_h, \text{class}_{\bar{h}}) \Big|_{\text{image}((\text{class}_h, \text{class}_{\bar{h}}))}$.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.1. Notation conventions and definitions

Si $\text{image}(h, h') = \text{image}(h) \times \text{image}(h')$ on dit que h et h' sont variation-independants ou séparés et on note $h \perp h'$ et $(h, h') \equiv \mathcal{E}$.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.1. Notation conventions and definitions

2.1.2. Modèle statistique général

- $\{(\Omega, \mathfrak{S}_\Omega, P)\}_{P \in \mathcal{P}}$, un modèle statistique
- $X : (\Omega, \mathfrak{S}_\Omega) \rightarrow (\mathcal{X}, \mathfrak{S}_\mathcal{X})$, la statistique observée
- $V : (\Omega, \mathfrak{S}_\Omega) \rightarrow (\mathcal{V}, \mathfrak{S}_\mathcal{V})$, latente, ou non. Non latente: $V < X$. Le **processus d'intérêt**
- La cible de l'inférence est d'estimer $\theta(P^V)$ ou de prédire $\theta(V)$. en résumé: $\theta : (\{P^V : P \in \mathcal{P}\} \times V(\Omega)) \rightarrow \cdot$.
- Soit $\bar{V}$ un complément (séparé ou pas) de V , latent ou non. La **nuisance**
- Soit $\mathcal{P}$ l'ensemble des distributions possibles de $\bar{V}$ une fois le processus ignoré.

Puisque $(V, \bar{V}) \sqsupseteq \Omega$, $\exists \mathbf{x} : (V(\Omega) \times \bar{V}(\Omega))$ est tel que $X = \mathbf{x}(V, \bar{V})$, $\mathbf{x}$ est donné par la relation: $\mathbf{x} = X[\cap_{V, \bar{V}}]$.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.2. Ignorer le processus de nuisance

2.2. Ignorer le processus de nuisance

2.2.1. Transformation du modèle

Pour $v \in V(\Omega)$, on définit l'ensemble mesurable

$$\Phi_{v,V,\bar{V}} = \bar{V}^{-1} \left(\bar{V} \left(V^{-1} (\{v\}) \right) \right).$$

Si $(V, \bar{V}) \equiv \Omega$, alors $\forall \bar{v} \in \bar{V}(\Omega)$, $\Phi_{\bar{v},\bar{V},V} = \Omega$, et

: $\omega \mapsto (\cap_{\bar{V},V}) (\bar{v}, V(\omega))$ est mesurable : $(\Omega, \mathfrak{S}_\Omega) \rightarrow (\Omega, \mathfrak{S}_\Omega)$.

Si $\forall P \in \mathcal{P}$, $\forall \bar{v} \in \bar{V}(\Omega)$, $P(\Phi_{\bar{v},\bar{V},V}) > 0$, on définit:

$$\mathcal{P}^* = \left\{ \left(\left\langle \bar{v} \mapsto P^{(V,\bar{v})|\Phi_{\bar{v},\bar{V},V}} \middle| P' \right\rangle_{(V,\bar{V})(\Omega),\bar{V}(\Omega)} \right)^{\cap_{V,\bar{V}}} \middle| (P, P') \in \mathcal{P} \times \mathcal{P}' \right\}.$$

Le modèle obtenu après avoir ignoré $\bar{V}$ est le modèle
 $(\text{codomain}(X), \mathfrak{S}_{\text{codomain}(X)}, \{(P^*)^X \mid P^* \in \mathcal{P}^*\})$.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.2. Ignorer le processus de nuisance

2.2.2. Transformation de la cible de l'inférence

Etant donné une cible θ , quelle est la cible dans le modèle correspondant θ^* dans le modèle $\mathcal{P}^*$? Dans certains cas, il y a une réponse naturelle

- ❶ S'il existe $Y : (\Omega, \mathfrak{S}_\Omega) \rightarrow \mathcal{Y}$ tel que $\theta : (\Omega \times \mathcal{P}) \rightarrow \mathcal{Y}, (\omega, P) \mapsto Y(\omega)$, un choix naturel est $\theta^* : \mathcal{P}^* \times \Omega \rightarrow \mathcal{Y}, (P^*, \omega) \mapsto (P^*, \omega) \mapsto Y(\omega)$.
- ❷ S'il existe $Y : (\Omega, \mathfrak{S}_\Omega) \rightarrow \mathcal{Y}$ tel que $\theta < (P, \omega \mapsto P^Y)$: $\theta = \theta' \circ (P \mapsto P^Y)$ et si $\{(P^*)^Y \mid P^* \in \mathcal{P}^*\} \subset \text{codomain}(\theta')$, alors on définit: $\theta^* : P^* \mapsto \theta'((P^*)^Y)$.
- ❸ Quand $\mathcal{P}^* \subset \mathcal{P}$, un choix naturel est $\theta^* = \theta \Big|_{\mathcal{P}^* \times \Omega}$.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.3. Transformation de l'ensemble des distributions a priori

2.3. Transformation de l'ensemble des distributions a priori

Après être passé de $\mathcal{P}$ à $\mathcal{P}^*$ après avoir ignoré un processus, il est difficile de définir naturellement un nouveau jeu de probabilités a priori $\mathcal{Q}^*$ parce que Q et Q^* ne sont pas des distributions définies sur le même ensemble. Quand chaque distribution est le produit d'un a priori sur le processus d'intérêt et d'une distribution sur le processus de nuisance, un modèle naturel peut être proposé:

- ❶ On définit l'ensemble $\tilde{\mathcal{Q}}$ de distributions à priori (pas forcément propre) sur $\mathcal{P}'$.

- ❷ Define $\mathcal{Q}^* =$

$$\left\{ \left(Q^{P \mapsto (\bar{v} \mapsto P^{(V, \bar{v}) | \Phi_{\bar{v}, \bar{V}, V})} \otimes Q' \right)^{(P'', P') \mapsto \left((\langle P'' | P' \rangle_{(V, \bar{V})(\Omega), \bar{V}(\Omega)})^{\cap_{V, \bar{V}}} \right)} : \right. \\ \left. Q \in \mathcal{Q}, Q' \in \mathcal{Q}' \right\}.$$

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.4. Inférences équivalentes

2.4. Inférences équivalentes

- Inférence basée sur la vraisemblance: L'inférence basée sur $\theta : \mathcal{P} \rightarrow$ est équivalente à l'inférence basée sur $\theta^* : \mathcal{P}^* \rightarrow$ lorsque $X = x$ if $\text{codomain}(\theta) = \text{codomain}(\theta^*)$ et
$$\left(\theta \mapsto \mathcal{L}_{\theta^*, X}^*(\theta; x) \right) = \left(\theta \mapsto \mathcal{L}_{\theta, X}(\theta; x) \right).$$
- Estimation fréquentiste:
L'inférence basée sur l'estimation de $\theta : \mathcal{P} \rightarrow$ par $\hat{\theta} : \Omega \rightarrow \text{codomain}(\theta)$ est équivalente à l'estimation de $\theta^* : \mathcal{P} \rightarrow$ par $\hat{\theta}$ si $\text{codomain}\theta = \text{codomain}\theta^*$ et $\forall \theta \in \text{codomain}(\theta),$
$$\{P^{\hat{\theta}} \mid P \in \theta^{-1}(\theta)\} = \{(P^*)^{\hat{\theta}} \mid (P^*) \in (\theta^*)^{-1}(\theta)\}.$$

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.4. Inférences équivalentes

- Inférence Bayésienne: L'inférence Bayésienne sur $\theta : (\mathcal{P} \times \Omega) \rightarrow$ avec l'ensemble de distributions à priori $\mathcal{Q}$ est équivalent à l'inférence basée sur la vraisemblance sur $\theta^* : (\mathcal{P}^* \times \Omega) \rightarrow$ avec le jeu de distributions a priori $\mathcal{Q}$ quand $X = x$ si $\{(Q^*)^{\theta^*|X=x} \mid Q^* \in \mathcal{Q}^*\} = \{Q^{\theta|X=x} \mid Q \in \mathcal{Q}\}$

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.5. Processus ignorable vs. processus informatif : définition et caractérisation

2.5. Processus ignorable vs. processus informatif : définition et caractérisation

2.5.1. Définition

Le processus $\bar{V}$ sera dit ignorable pour un type particulier d'inférence si pour $\mathcal{P}' = \{\text{Dirac}_{\bar{v}} : \bar{v} \in \bar{V}(\Omega)\}$, l'inférence basée sur $(\mathcal{P}, \theta, \mathcal{L}, \mathcal{Q}, \dots)$ est équivalente à l'inférence basée sur $(\mathcal{P}^*, \theta^*, \mathcal{L}^*, \mathcal{Q}^*, \dots)$. Dans le cas où les deux ne sont pas équivalents, le processus $\bar{V}$ est dit informatif.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.5. Processus ignorable vs. processus informatif : définition et caractérisation

De façon regrettable, “ V est ignorable” dans ce sens ne signifie pas que “ V is ignorable” dans le sens habituel de l’ignorabilité en statistiques. Une statistique ignorable est toute statistique indépendante d’au moins une variable exhaustive (Schervish 1995, p. 142, Exercise 33).

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.5. Processus ignorable vs. processus informatif : définition et caractérisation

Exemple:

$Y \sim \text{Normal}(0, \sigma^2)^{\otimes N}$, et T est la sélection associée avec $R = (\arg \min(Y_k \mid k \in \{1, \dots, N\}), \arg \max(Y_k \mid k \in \{1, \dots, N\}))$. On observe T et $T[Y]$ Ici: $T[Y] = (\min(Y), \max(Y))$ et $T[Y]$ est indépendant de T . Etant donné $R = (1, 5)$, T est la fonction $t : y \mapsto (y_1, y_5)$, et $P^{t(Y)} = P^{(Y_1, Y_5)} \neq P^{T[Y] \mid T=t} = P^{\min(Y), \max(Y)}$. Donc la sélection est non ignorable et informative, bien que $P_{\theta, \bar{\theta}}^T$ soit indépendant de θ , et inconditionnellement à Y , T est le sondage aléatoire simple de 2 parmi N . donc la variable T est libre, et ignorable puis que indépendante de la statistique exhaustive $T[Y]$.

2. Modèle générique pour les sondages et sélection informative dans la littérature scientifique

2.6. Au hasard (at random)

2.6. Au hasard (at random)

La condition “At random” est la condition ou

$$\mathcal{P} = \{P^V \otimes P^{\bar{V}} : P \in \mathcal{P}\}.$$

3

Conditionnellement au
modèle, quelle est la meilleure
inférence ?

3. Conditionnellement au modèle, quelle est la meilleure inférence ?

Pour résumer:

- Rubin, 1976, donne une définition rigoureuse et généralisable d'un processus de non réponse ignorable.
- On se retrouve à se poser la question de la meilleure inférence en présence d'un processus de nuisance.
- Quelles sont les meilleures prédictions pour un modèle donné ?

3. Conditionnellement au modèle, quelle est la meilleure inférence ?

Premier écueil: ne pas spécifier le modèle ! Une inférence est optimale pour un modèle donné.

3. Conditionnellement au modèle, quelle est la meilleure inférence ?

3.1. Cas fréquentiste



Basu, Debabrata. 1977.

On the Elimination of Nuisance Parameters.

Journal of the American Statistical Association, 72(358),
355—366.

Pour l'inférence classique, il n'y a pas de définition totalement satisfaisante d'une statistique exhaustive en présence d'un paramètre de nuisance.

3. Conditionnellement au modèle, quelle est la meilleure inférence ?

3.1. Cas fréquentiste

Marginaliser

$$\mathcal{L}_{S,Y[S]}(\mathbf{x}, \mathbf{y}; \theta, \xi) = f_{Y[\mathbf{x}];\theta}(\mathbf{y}) \times f_{S|Y[\mathbf{x}];\theta,\bar{\theta}}(\mathbf{x} | \mathbf{y})$$

ou

conditionner

$$\mathcal{L}_{Y[S]|S}(\mathbf{y} | \mathbf{x}; \theta, \bar{\theta}) = f_{Y[\mathbf{x}];\theta}(\mathbf{y}) \times \frac{f_{S|Y[\mathbf{x}];\theta,\bar{\theta}}(\mathbf{x} | \mathbf{y})}{f_{S;\theta,\bar{\theta}}(\mathbf{x})}$$

?

- Conflit Fisher - Barnard

3. Conditionnellement au modèle, quelle est la meilleure inférence ?

3.2. Cas fréquentiste - modèle fixe de population, inférence basée sur le plan

Pas de résultat d'optimalité.

4

Quel est le meilleur modèle ?

4. Quel est le meilleur modèle ?

4.1. Cas de l'inférence basée sur le plan

Dans un institut de sondage, en l'absence de non réponse, l'estimateur d'Horvitz Thompson est:

$$\sum_{k \in U} \frac{Y_k}{\pi_k} \mathbb{1}_S(k).$$

et l'estimateur de la variance est:

$$\sum_{k \in U} Y_k Y_\ell (\pi_{k,\ell})^{-1} (\pi_{k,\ell} - \pi_k \pi_\ell) \mathbb{1}_S(k) \mathbb{1}_S(\ell).$$

avec $\pi_{k,\ell} = E[\mathbb{1}_S(k) \mathbb{1}_S(\ell)]$

L'estimation de la précision de notre mesure est basée sur ce que d'autres échantillons auraient pu nous donner:

$$V[\hat{t}_y(S)] = \sum_{s \in S(\Omega)} (\hat{t}_y(s) - t_y)^2 P(S = s).$$

4. Quel est le meilleur modèle ?

4.1. Cas de l'inférence basée sur le plan

Un directeur d'institut a 3 sondeurs préférés, plus ou moins matinaux.

- Le premier interroge les gens de son village
- Le deuxième fait du sondage aléatoire simple
- Le troisième un plan stratifié avec allocation constante.

Il décide que le premier à lui apporter un café sera en charge du prochain plan de sondage.

Quelles sont les probabilités d'inclusion et les probabilités de double inclusion ?

4. Quel est le meilleur modèle ?

4.1. Cas de l'inférence basée sur le plan

Eléphants de Basu : Pour un modèle fixe de population,

- le refus de sortir du modèle fixe de population empêche de formuler le lien entre information auxiliaire et variable d'intérêt.

4. Quel est le meilleur modèle ?

4.1. Cas de l'inférence basée sur le plan

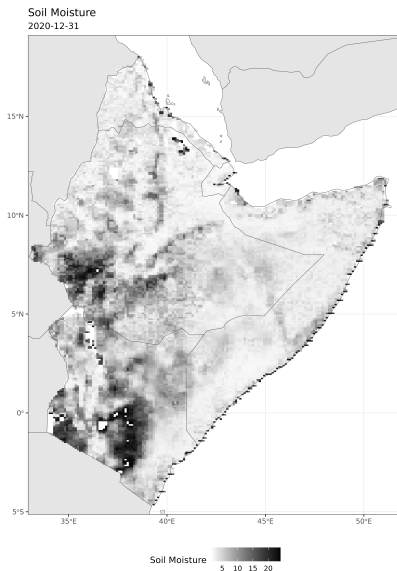
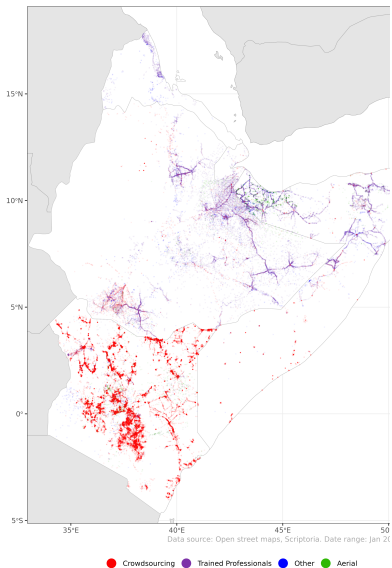
Le choix optimal (déjà appliqué):

Hors enquêtes de type enquête nationales, ou l'estimateur est suffisamment concentré autour du paramètre d'intérêt (petits domaines, non réponse, données participatives) :

- dépasser le dogme de l'inférence basée sur le plan
- modéliser la sélection conditionnellement à la variable d'intérêt.
- essayer de proposer une inférence raisonnable

4. Quel est le meilleur modèle ?

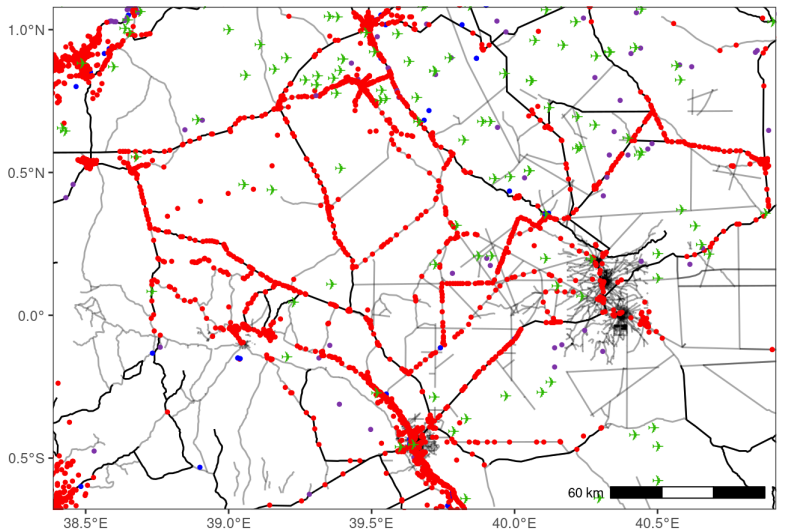
4.2. Un exemple



4. Quel est le meilleur modèle ?

4.2. Un exemple

Dans la littérature on trouve l'utilisation des images à haute résolution pour expliquer pourquoi un criquet était à un endroit donné et pas 10 mètres plus loin.



● Crowdsourcing ● Trained Professionals ● Other ✈ Aerial — Main road — Secondary ro.

4. Quel est le meilleur modèle ?

4.2. Un exemple

- Données participatives
- Erreur de mesure (spatiale, temporelle, sur la nature des observations)
- Campagnes de mesure (stratification temporelle)
- Non uniformité spatiale
- Pas d'information sur la planification des campagnes

4. Quel est le meilleur modèle ?

4.2. Un exemple

Nous sommes en présence de deux processus:

- Le processus de nuisance c'est la rédaction d'un rapport.
- Le processus d'intérêt c'est la présence de criquets.

Les deux sont dépendants des variables environnementales.

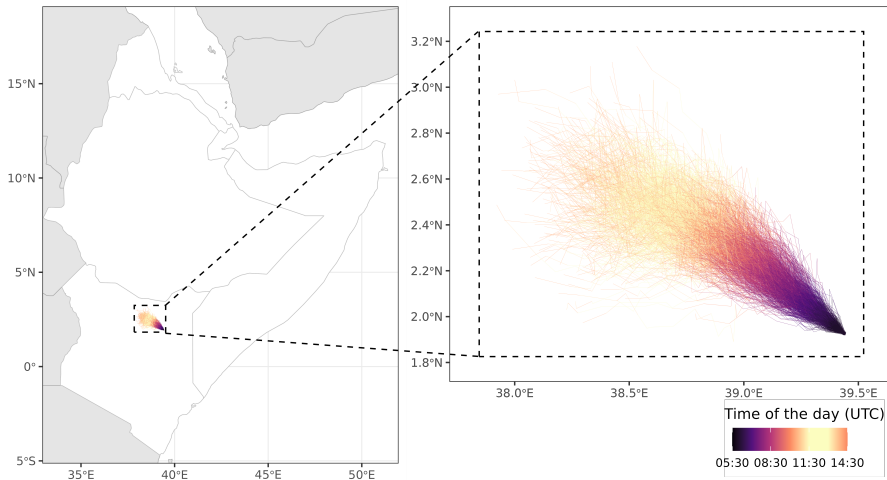
4. Quel est le meilleur modèle ?

4.2. Un exemple

Objectif: prédiction de court terme sur la direction des essais.

4. Quel est le meilleur modèle ?

4.2. Un exemple



4. Quel est le meilleur modèle ?

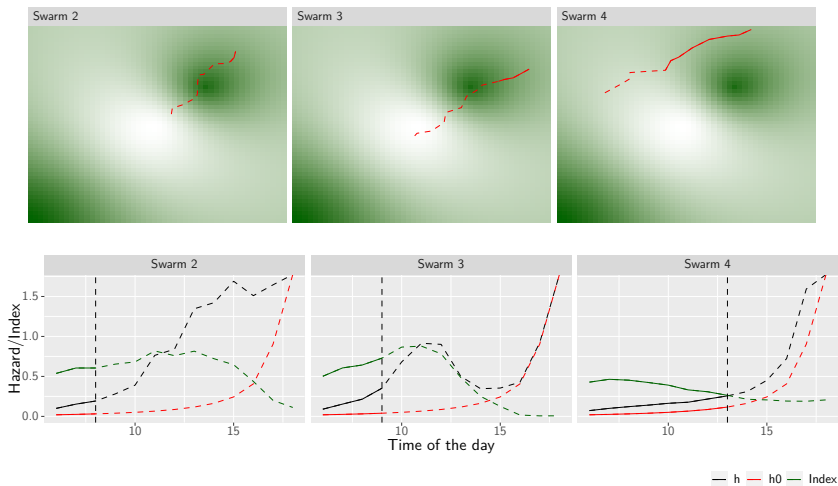
4.2. Un exemple

- Modèle de survie pour la distribution des heures d'atterrissage
- Modélise l'intensité de surveillance comme une fonction de la présence d'une route et de la densité de population et du pays pour chaque type de surveillance

4. Quel est le meilleur modèle ?

4.2. Un exemple

Figure 4.1: Single swarm potential and realised trajectories



4. Quel est le meilleur modèle ?

4.2. Un exemple

Principe de l'inférence (Sequential Monte Carlo ABC):

- Tirer les paramètres dans la loi a priori (initialisation)
- puis récursivement :
 - Simuler les observations et les vols de criquets
 - Calculer une distance entre observations et observations simulées
 - Conserver 10% des paramètres donnant les plus petites distances
 - Retirer les paramètres dans le voisinage des paramètres sélectionnés

Sont prédits à la fois:

- Le nombre de criquets
- Le nombre de signalements

Est obtenu la loi jointe a posteriori des paramètres de nuisance et d'intérêt.

5

Sélection optimale

Soit P' un ensemble de lois pour le processus de nuisance (plan à taille fixe n) par exemple.

Un critère d'optimalité possible serait :

$$\arg \min_{P' \in \mathcal{P}'} \int_{P \in \mathcal{P}} \mathbb{E} \left[\left(\hat{\theta}(X) - \theta(V, P) \right)^2 \right] dQ(P)$$

avec $\hat{\theta}(X) = E[\theta(V, P) \mid X]$.

Problèmes :

- ce critère n'intègre pas le cout de calcul de $\hat{\theta}(X)$, qui n'est pas calculable quelque soit P'
- l'optimum est incalculable.

Exemple :

- ❶ Données méta génomiques
- ❷ Echantillonnage adaptatif



6

Conclusion

6. Conclusion

Merci de votre attention