

L'INFÉRENCE STATISTIQUE, ESTIMATION ET SONDAGES

Michel Henry,
IREM de Franche-Comté

Sommaire

I	Principes de l'inférence statistique	39
1	Populations statistiques et échantillons	39
2	Modèle probabiliste de l'échantillonnage	41
3	Modélisation des paramètres standards	42
4	Propriétés des statistiques $\bar{X}$ et V	43
II	Estimation	45
1	Principe de l'estimation	45
2	Qualités d'une estimation ponctuelle	46
3	Estimation par intervalle de confiance	47

I Principes de l'inférence statistique

1 Populations statistiques et échantillons

On désire mieux connaître une **population** P (d'êtres vivants ou d'objets, comme par exemple une production de marchandises, ou tout autre ensemble relevant d'études statistiques) et notamment certaines de ses caractéristiques.

L'investigation exhaustive n'est généralement pas possible : coûts exorbitants, inaccessibilité de l'ensemble P , temps disponible pour obtenir les résultats de l'étude. . .

On peut penser par exemple à la différence de coût et de méthode entre un recensement de la population française (périodiquement nécessaire) et un sondage dont la pratique dépend d'hypothèses générales basées sur un recensement antérieur.

L'étude de la population conduit à définir des **caractères** qui attribuent aux éléments de P certaines qualités ou valeurs qui résument les caractéristiques étudiées. L'ensemble des valeurs prises par un caractère χ fournit une **série statistique** qui peut être répartie en **classes** et représentée par divers **diagrammes**.

Quand le caractère χ est **quantitatif**, la série peut à son tour être résumée par certains **paramètres**, notamment la **moyenne** μ de χ et son **écart-type** σ . Ces valeurs

sont en général inconnues lorsque l'on entreprend une étude statistique. L'un de ses objets est d'en donner des approximations suffisantes pour prendre des décisions.

Quand χ est **qualitatif**, on s'intéresse aux proportions dans P de ses diverses modalités. Par exemple, on désire connaître le poids p de l'une d'entre elles (situation des sondages).

Le principe de l'inférence statistique est d'obtenir des informations sur la population P (population « mère ») à partir de la connaissance d'un échantillon E .

Précisons d'abord ce que l'on entend par « **échantillon** ».

Dans la réalité statistique, un n -échantillon est un ensemble de n objets prélevés dans la population. Quand le prélèvement se fait au hasard, on obtient un **échantillon aléatoire**.

Quand la population P est partagée en strates, dont les poids respectifs sont connus (à partir d'un recensement), on peut décider *a priori* de composer l'échantillon proportionnellement aux poids des strates dans la population. On dit alors qu'on a un « **échantillon représentatif** » de P . Ce n'est pas toujours le meilleur quant aux résultats que l'on veut en tirer^(a).

Pour appliquer certains théorèmes de probabilité, il convient de faire l'hypothèse que les éléments de l'échantillon E sont prélevés indépendamment les uns des autres. Pour cela, afin que les prélèvements successifs ne modifient pas la composition de la population, il faudrait les effectuer avec remises, ce qui n'est pas toujours possible. Dans la pratique, les échantillons sont plutôt prélevés sans remises. Lorsque la population est vaste par rapport à la taille n de l'échantillon (au moins 1000 fois plus grande), cet inconvénient est mineur, entachant les probabilités en jeu d'une erreur négligeable.

En inférence statistique, quand on s'intéresse à un caractère quantitatif χ on peut **estimer** (évaluation statistique) les valeurs des paramètres qui résument la distribution du caractère χ sur la population P , à partir des valeurs de χ observées sur l'échantillon E . On peut aussi estimer la valeur de la proportion p des éléments de P qui appartiennent à une certaine modalité d'un caractère qualitatif χ , à partir de la fréquence f de cette modalité dans E . C'est notamment la situation des sondages aléatoires simples.

On peut aussi émettre des **hypothèses** concernant P et tester la validité de ces hypothèses à partir des renseignements que fournit l'échantillon (tests statistiques).

Quand l'échantillon est aléatoire, **cette inférence est en partie déterminée par le hasard** du prélèvement. Ainsi, toute déclaration à propos de la population, issue de l'observation d'un échantillon, est entachée d'une certaine **risque** (probabilité) de se

^(a)Car les valeurs du caractère sur certaines strates peuvent présenter une trop grande dispersion, ce qui peut faire surestimer le poids de la strate et augmenter la dispersion résultante. Or cette dispersion intervient pour déterminer un encadrement de la valeur pour la population P du paramètre étudié.

tromper. Le problème de l'inférence statistique est de pouvoir donner des résultats suffisamment précis avec un risque minimisé, deux contraintes qui varient en sens contraire. Pour améliorer à la fois la **précision** des résultats et la **fiabilité** des déclarations relatives à ces résultats, le statisticien ne peut qu'agrandir la taille de son échantillon, ce qui coûte plus cher.

2 Modèle probabiliste de l'échantillonnage

Nous allons maintenant construire le modèle mathématique général pour décrire une situation d'échantillonnage en précisant le problème :

Soit E un échantillon aléatoire de taille n , extrait de la population P .

Le caractère χ étudié a pour moyenne μ et pour écart type σ dans P .

Les éléments de P qui donnent à χ une certaine modalité A sont en proportion p .

La moyenne de χ sur E est m_n et son écart type est s_n .

La fréquence de la modalité A dans E est f_n .

Quel lien y a-t-il entre les valeurs inconnues μ , σ , p et les valeurs observées m_n , s_n et f_n ?

Dans le modèle probabiliste, on représente le prélèvement au hasard d'un élément de la population P par un élément ω d'un **ensemble référentiel** Ω , que l'on supposera fini. Ω symbolise donc à la fois les éléments de P et le fait qu'ils peuvent être issus d'un tirage aléatoire dans lequel les éléments de P ont la même « chance » d'être choisis.

Pour représenter un caractère quantitatif χ étudié (pour simplifier, on le supposera de dimension 1), on considère une **variable aléatoire** (v.a.) X_0 , définie sur Ω , prenant pour chaque **éventualité** ω , la valeur du caractère χ pour l'élément de P dont ω représente le tirage au hasard. Pour un caractère qualitatif, X_0 prendra les valeurs 1 ou 0 suivant que pour cet élément, χ prend ou non la modalité A .

La variable X_0 est censée décrire la répartition des valeurs de χ , avec leurs poids respectifs. Sa **loi** (ensemble des probabilités des événements associés aux valeurs possibles de χ) modélise cette répartition.

Dans le cas d'un caractère quantitatif, on a noté μ la moyenne de χ sur la population P et σ^2 sa variance. La loi de X_0 est inconnue, mais on fait l'hypothèse que ces valeurs μ et σ^2 sont les valeurs de l'**espérance mathématique** et de la **variance** de cette loi : $E(X_0) = \mu$ et $\text{Var}(X_0) = \sigma^2$.

Dans le cas d'une proportion p , la loi de X_0 est donnée par $\mathbb{P}(X_0 = 1) = p$ (probabilité de tirer au hasard de P un élément présentant la modalité A) et on a :

$$E(X_0) = p \quad \text{et} \quad \text{Var}(X_0) = p(1 - p).$$

La réalisation de l'échantillon E fournit donc un n -uplet de valeurs observées de $\chi : x = (x_1, \dots, x_i, \dots, x_n)$, interprétées comme images d'événements $\{\omega_i \in \Omega\}$ par l'application répétée n fois de X_0 . On considère que ces x_i sont les réalisations (observations) de n v.a. X_i qui représentent cette réplique successive de X_0 .

On se place donc dans une hypothèse de travail : les éléments de E sont prélevés au hasard et indépendamment les uns des autres (notamment la taille de P est assez grande par rapport à celle de E).

Cette hypothèse de travail se traduit par une hypothèse de modèle : les X_i sont des variables aléatoires indépendantes (au sens probabiliste) et de même loi^(b) que X_0 . Notamment, on a pour tous les i : $E(X_i) = E(X_0)$ et $\text{Var}(X_i) = \text{Var}(X_0)$.

L'échantillonnage est donc représenté par un **vecteur aléatoire** $X = (X_1, \dots, X_i, \dots, X_n)$ vérifiant cette hypothèse de modèle. Le résultat effectif de cet échantillonnage est donc un échantillon E , représenté par les valeurs observées $(x_1, \dots, x_i, \dots, x_n)$. X est encore appelé un « **échantillon de la v.a. X_0** » et X_0 est appelée « **variable parente** » de l'échantillon X .

Par définition, dans le modèle de la statistique :

Un échantillon est un n -uplet de variables aléatoires, définies sur le même Ω , indépendantes et de même loi.

3 Modélisation des paramètres standards

Pour un caractère quantitatif, la moyenne m_n de χ sur l'échantillon E est représentée dans le modèle probabiliste par la moyenne arithmétique $\bar{x} = \frac{1}{n} \sum x_i$ (les probabilistes ont l'habitude de désigner cette moyenne par le symbole $\bar{x}$). $\bar{x}$ est donc la valeur observée d'une variable aléatoire $\bar{X} = \frac{1}{n} \sum X_i$ définie sur Ω^n , par l'intermédiaire de X .

La variance sur E du caractère χ est aussi un paramètre important. Sa valeur est donnée par : $s_n^2 = \frac{1}{n} \sum (x_i - \bar{x})^2$. C'est la valeur observée^(c) de la variable aléatoire $V = \frac{1}{n} \sum (X_i - \bar{X})^2$ également définie sur Ω^n , par l'intermédiaire de X .

Pour l'étude d'une proportion, $\bar{x}$ représente aussi la fréquence observée f_n de la modalité A dans l'échantillon E , puisque le $\sum$ contient autant de nombres 1 qu'il y a d'éléments dans E de modalité A . Comme $\text{Var}(X_0) = p(1-p)$, p est le seul paramètre inconnu pour caractériser la population du point de vue de la modalité A . L'introduction de la variable V est alors inutile.

^(b) X_i est de même loi que X_0 si pour tout x , $\mathbb{P}(X_i = x) = \mathbb{P}(X_0 = x)$. Les X_i sont indépendantes si pour tout n -uplet $(x_1, \dots, x_i, \dots, x_n)$, on a

$$\mathbb{P}[(X_1, \dots, X_i, \dots, X_n) = (x_1, \dots, x_i, \dots, x_n)] = \mathbb{P}[X_1 = x_1] \cdot \mathbb{P}[X_2 = x_2] \dots \mathbb{P}[X_n = x_n].$$

^(c) Rappelons que les valeurs $\bar{x}$ et s_n^2 sont issues d'un échantillonnage aléatoire et ne coïncident pas avec la moyenne et la variance inconnues μ et σ^2 de χ sur la population P . Il en est de même pour f_n et p .

Ces variables $\bar{X}$ et V sont aussi appelées des « résumés statistiques » (on dit plus brièvement « statistiques »^(d)), car leurs réalisations combinent les valeurs de χ sur E .

L'étude des lois des variables $\bar{X}$ et V , en fonction éventuellement d'hypothèses supplémentaires sur la répartition de χ dans P , est une partie importante de la statistique inférentielle. Elles jouissent de propriétés générales obtenues à partir de certains théorèmes puissants de la théorie des probabilités. Nous utiliserons les suivants :

Espérance mathématique

1. L'espérance mathématique est une forme linéaire sur l'espace des v.a. définies sur le même Ω . Si Ω est un ensemble fini sur lequel une v.a. Y prend les valeurs y_k et dont la loi est donnée par les probabilités $p_k = \mathbb{P}(Y = y_k)$, on a $E(Y) = \sum p_k y_k$.
2. Si Y et Z , définies sur Ω , sont deux variables indépendantes (tout événement lié à l'une est indépendant de tout événement lié à l'autre), on a : $E(Y.Z) = E(Y).E(Z)$.

Variance

1. Avec les hypothèses et notations précédentes, la variance de Y est définie par $\text{Var}(Y) = \sum p_k (y_k - E(Y))^2$. On a pour α réel : $\text{Var}(\alpha Y) = \alpha^2 \text{Var}(Y)$.
2. Si Y et Z sont indépendantes, on a : $\text{Var}(Y + Z) = \text{Var}(Y) + \text{Var}(Z)$.

Inégalité de BIENAYMÉ-TCHEBYCHEV

On montre assez facilement l'inégalité suivante : $\mathbb{P}(|Y - E(Y)| > \varepsilon) \leq \frac{\text{Var}(Y)}{\varepsilon^2}$.

Cette inégalité permet de majorer la probabilité de se tromper (le risque que l'on prend) quand, à partir d'une observation y de Y , on affirme que la valeur moyenne $E(Y)$ est dans l'intervalle $]y - \varepsilon, y + \varepsilon[$. Elle permet donc d'avoir une idée de la fiabilité de cette affirmation. Notons que la majoration de cette probabilité varie comme la dispersion (variance) de Y et est en $\frac{1}{\varepsilon^2}$.

4 Propriétés des statistiques $\bar{X}$ et V

- $\bar{X}$, **moyenne de l'échantillon**, a des propriétés sympathiques. La linéarité de l'espérance mathématique donne immédiatement :

$$E(\bar{X}) = E\left(\frac{1}{n} \sum_{i=1}^n X_i\right) = \frac{1}{n} \sum_{i=1}^n E(X_i) = \frac{1}{n} n.E(X_0) = E(X_0).$$

D'où le résultat :

$$E(\bar{X}) = \begin{cases} \mu & \text{pour un caractère quantitatif} \\ p & \text{pour un caractère qualitatif} \end{cases}$$

^(d)Une statistique est donc une application réelle définie sur Ω^n , composée de X et d'une application de $\mathbb{R}^n$ dans $\mathbb{R}$.

- Les propriétés de la variance, compte tenu de l'indépendance des X_i , donnent aussi :

$$\text{Var}(\bar{X}) = \text{Var}\left(\frac{1}{n} \sum_{i=1}^n X_i\right) = \frac{1}{n^2} \sum_{i=1}^n \text{Var}(X_i) = \frac{\text{Var}(X_0)}{n}.$$

D'où :

$$\text{Var}(\bar{X}) = \begin{cases} \frac{\sigma^2}{n} & \text{pour un caractère quantitatif} \\ \frac{p(1-p)}{n} & \text{pour un caractère qualitatif} \end{cases}$$

- **Application à la précision de mesures physiques.**

Anticipons un peu sur le paragraphe consacré à l'estimation par intervalle de confiance, pour montrer comment ces résultats justifient une intuition forte : sans changer d'appareil de mesure, on peut améliorer notablement sa **précision** (marge d'erreur sur le résultat annoncé) en faisant la moyenne de plusieurs mesures indépendantes.

Supposons que la variable parente X_0 représente la mesure d'une grandeur (physique par exemple) inconnue μ . L'inégalité de BIENAYMÉ-TCHEBYCHEV, appliquée à X_0 avec $\varepsilon = \varepsilon' \sigma$ donne :

$$\mathbb{P}[X_0 - \varepsilon' \sigma \leq E(X_0) \leq X_0 + \varepsilon' \sigma] \geq 1 - \frac{1}{\varepsilon'^2}.$$

L'intervalle $[x_0 - \varepsilon' \sigma; x_0 + \varepsilon' \sigma]$ obtenu pour encadrer $\mu = E(X_0)$ à partir d'une simple observation x_0 de X_0 , a une longueur proportionnelle à l'écart type σ de X_0 , pour une **fiabilité** $1 - \frac{1}{\varepsilon'^2}$ donnée (minorant de la probabilité d'annoncer un bon encadrement).

Appliquée à $\bar{X}$, avec $\varepsilon = \frac{\varepsilon' \sigma}{\sqrt{n}}$, on a : $\mathbb{P}\left[\bar{X} - \frac{\varepsilon' \sigma}{\sqrt{n}} \leq E(\bar{X}) \leq \bar{X} + \frac{\varepsilon' \sigma}{\sqrt{n}}\right] \geq 1 - \frac{1}{\varepsilon'^2}.$

La « **fourchette** » $\left[\bar{x} - \frac{\varepsilon' \sigma}{\sqrt{n}}; \bar{x} + \frac{\varepsilon' \sigma}{\sqrt{n}}\right]$ proposée cette fois-ci pour l'encadrement de cette valeur inconnue $\mu = E(\bar{X})$, montre que l'utilisation de $\bar{X}$ plutôt que X_0 divise la marge d'erreur par $\sqrt{n}$. On voit aussi que, pour un n donné, on améliore la précision (on resserre la fourchette en diminuant ε') en consentant une perte de fiabilité $\left(\text{en } 1 - \frac{1}{\varepsilon'^2}\right)$: on ne peut pas avoir le beurre et l'argent du beurre ! Même remarque pour l'estimation d'une proportion.

- **Espérance de V**

On est dans le cas où χ est un caractère quantitatif.

La variance $V = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2$ de l'échantillon est moins agréable que sa moyenne. Pour l'étudier, il est commode d'introduire la variance de χ restreinte à l'échantillon. Pour cela, on considère la statistique notée Σ^2 , définie par :

$$\Sigma^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \mu)^2.$$

On a la relation^(e) : $V = \Sigma^2 - (\bar{X} - \mu)^2$.

Pour calculer $E(V)$, on utilise les propriétés de linéarité de l'espérance mathématique :

$$\begin{aligned} E(V) &= E(\Sigma^2) - E\left[(\bar{X} - \mu)^2\right] \\ &= \frac{1}{n} \sum_{i=1}^n E[(X_i - \mu)^2] - E\left[(\bar{X} - \mu)^2\right] \\ &= \text{Var}(X_0) - \text{Var}(\bar{X}) \end{aligned}$$

D'où :

$$E(\Sigma^2) = \sigma^2 \text{ et } E(V) = \frac{n-1}{n} \sigma^2$$

II Estimation

1 Principe de l'estimation

L'estimation d'un paramètre résumant les valeurs d'un caractère dans une population, consiste à donner pour la valeur τ prise par ce paramètre pour la population, une valeur approchée t_n , calculée à partir d'un échantillon, avec autant de précision que possible.

On utilise pour cela une statistique T appropriée, en retenant pour τ la valeur $t_n = T(x)$, où x est l'observation de l'échantillon X .

Par exemple, pour estimer une moyenne μ ou une proportion p , on utilisera $\bar{X}$. Pour estimer une dispersion, on peut penser à la statistique V (qui n'est pas la meilleure !). Ces variables prennent alors le nom d'« **estimateurs** ».

Pour une **estimation ponctuelle**, on se contente des valeurs observées sur l'échantillon : $\bar{x}$ valeur de $\bar{X}$, m_n ou f_n suivant les cas pour estimer les valeurs théoriques μ (moyenne d'un caractère quantitatif) ou p (fréquence d'un caractère qualitatif) d'une part, s_n^2 valeur observée de V pour estimer la dispersion σ^2 du caractère dans la population d'autre part.

Le contrôle de la marge d'erreur renvoie à une **estimation par intervalle**. Par exemple, comment majorer la probabilité de se tromper en donnant pour une moyenne inconnue μ un encadrement issu de la valeur observée $\bar{x}$ sur l'échantillon :

$$\mathbb{P}\left(\left|\bar{X} - \mu\right| < \varepsilon\right) \geq 1 - \alpha ?$$

^(e)On trouve le lien entre V et Σ^2 en ramenant dans V les variables X_i et $\bar{X}$ à leur espérance m , en écrivant : $X_i - \bar{X} = (X_i - m) + (m - \bar{X})$ et en développant le carré :

$$V = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{1}{n} \sum_{i=1}^n \left[(X_i - m)^2 + 2(X_i - m)(m - \bar{X}) + (\bar{X} - m)^2 \right],$$

d'où $V = \Sigma^2 + \frac{2}{n} (n\bar{X} - nm) (m - \bar{X}) + (\bar{X} - m)^2 = \Sigma^2 - (\bar{X} - m)^2$.

Un niveau de confiance $1 - \alpha$ étant donné (0,9 ou 0,95 suivant les degrés de fiabilité que l'on souhaite), si l'on sait calculer cette probabilité, on peut obtenir une valeur minimale pour ε (le 1/2 écartement de la **fourchette**), telle qu'avec une probabilité meilleure que $1 - \alpha$ on puisse affirmer que μ est dans l'intervalle $\left] \bar{x} - \varepsilon ; \bar{x} + \varepsilon \right[$.

La détermination de cet ε dans diverses situations concrètes est le problème du calcul des **intervalles de confiance**.

2 Qualités d'une estimation ponctuelle

Quand on donne une valeur expérimentale, issue de l'observation d'un échantillon aléatoire pour un paramètre inconnu, on peut se demander :

En quoi cette estimation ponctuelle est-elle une « bonne » estimation ?

On aimerait que cette inférence statistique jouisse des deux propriétés suivantes :

1. Si, pour chaque échantillon, chaque valeur observée t peut être différente de la valeur τ à estimer (cela dépend du hasard de l'échantillonnage), dans l'ensemble des échantillons de taille n possibles, on aimerait que ces valeurs se répartissent autour de τ de telle sorte qu'en moyenne, elles donnent τ . Autrement dit, on souhaite que $E(T) = \tau$. On dit alors que l'estimateur T est « **sans biais** ».
2. On aimerait aussi que la précision et la fiabilité de l'estimation s'améliorent en augmentant la taille de l'échantillon (cf. l'exemple ci-dessus des mesures physiques). On dit alors que T est un estimateur **convergent** : plus on prend de grands échantillons, plus les valeurs observées t sont proches de τ , ou du moins la probabilité qu'elles s'en rapprochent tend vers 1.

Pour vérifier ces propriétés, on doit faire appel aux résultats théoriques en probabilités que l'on a rappelés et aux hypothèses de modèle. Voyons de plus près ce qu'il en est pour les deux estimateurs $\bar{X}$ et V introduits.

Estimation d'une moyenne ou d'une proportion

Dans le cas d'un caractère χ quantitatif, pour estimer la **moyenne** μ de χ prenons l'estimateur $\bar{X}$. Comme $E(\bar{X}) = \mu$, $\bar{X}$ est un estimateur sans biais de μ .

De plus, $\text{Var}(\bar{X}) = \frac{\sigma^2}{n}$. L'inégalité de BIENAYMÉ-TCHEBYCHEV s'écrit :

$$\mathbb{P}\left(\left|\bar{X} - m\right| < \varepsilon\right) \geq 1 - \frac{\sigma^2}{n\varepsilon^2}$$

ce qui montre que $\bar{X}$ est un estimateur convergent.

De même, pour estimer une **proportion** p , l'estimateur $\bar{X}$ est performant : comme $E(\bar{X}) = p$, $\bar{X}$ est encore un estimateur sans biais, et $\text{Var}(\bar{X}) = \frac{p(1-p)}{n}$ montre qu'il est convergent.

On retrouve un résultat attendu : pour estimer la probabilité d'un événement A , on peut prendre la **fréquence** des réalisations de A lors de la répétition d'un (assez) grand nombre d'expériences dont les issues constituent l'échantillon E .

Ce résultat ne démontre pas que la loi des grands nombres est une nécessité naturelle, il ne fait que donner une confirmation de l'efficacité du modèle probabiliste qui permet notamment de démontrer l'inégalité de BIENAYMÉ-TCHEBYCHEV.

Historiquement ce résultat a été le premier théorème important du calcul des probabilités, démontré par Jacques BERNOULLI dans *Ars Conjectandi* (1713), conséquence d'un théorème connu aujourd'hui sous le nom de « loi faible des grands nombres ». Le théorème de BERNOULLI montre que la fréquence observée est un estimateur sans biais convergeant pour la probabilité d'un événement.

Mais il ne permet pas évaluer assez finement l'erreur commise par cette estimation ponctuelle, c'est à dire de donner un encadrement de confiance pour cette probabilité. Un tel encadrement (intervalle de confiance) repose sur un théorème dû à MOIVRE et LAPLACE (démontré en 1814), cas particulier d'un résultat puissant connu aujourd'hui sous le nom de « théorème limite central ». C'est ce théorème qui fait de la statistique inférentielle un outil efficace.

Estimation d'une variance

Par contre, $E(V) = \frac{n-1}{n} \sigma^2$, V est donc un estimateur biaisé de σ^2 . Mais $E\left(\frac{n}{n-1} V\right) = \sigma^2$, et par conséquent la statistique $S^2 = \frac{n}{n-1} V$, qui s'écrit aussi :

$$S^2 = \frac{1}{n-1} \sum (X_i - \bar{X})^2$$

est un estimateur sans biais de σ^2 .

Habituellement, on désigne par s^2 ou σ_{n-1}^2 la valeur observée de S^2 sur l'échantillon et on l'appelle la **variance estimée** de la population (les calculettes statistiques présentent les deux valeurs σ_n et σ_{n-1} , ça fait plus riche). $s^2 = \sigma_{n-1}^2$ est donc une estimation ponctuelle sans biais de la variance σ^2 du caractère χ dans la population. On démontre que S^2 est aussi un estimateur convergent de σ^2 .

Remarque : $E(S^2) = \sigma^2$ n'entraîne pas $E(S) = \sigma$, S n'est pas un estimateur sans biais de σ !

3 Estimation par intervalle de confiance

a) Niveau de confiance

On veut estimer la valeur d'un paramètre τ relatif à un caractère χ défini sur une population P . Une estimation ponctuelle à partir d'un échantillon ne renseigne pas sur la précision de l'approximation de τ . On voudrait donc obtenir un « intervalle aléatoire », pas trop grand,

à partir de l'échantillon prélevé, tel que la probabilité qu'il contienne τ soit acceptable.

Cette probabilité sera appelée « **niveau de confiance** » de l'estimation, on la désigne par $1 - \alpha$. Le nombre α est le risque que l'on prend de se tromper en affirmant que τ est bien dans l'intervalle proposé.

Pour préciser cela, prenons un niveau de confiance de 90 %. À chaque échantillon correspond la valeur observée t de l'estimateur T utilisé.

On considère l'intervalle centré en t : $]t - \varepsilon ; t + \varepsilon[$, où ε est choisi de sorte qu'en moyenne, pour 9 échantillons sur 10, τ soit dans $]t - \varepsilon ; t + \varepsilon[$.

Autrement dit, on désire trouver ε tel que $\mathbb{P}(\tau \in]T - \varepsilon ; T + \varepsilon[) \geq 0,9$. On a rencontré cette situation dans le cas où τ est l'espérance mathématique de la variable parente. On a vu que l'inégalité de BIENAYMÉ-TCHEBYCHEV donne alors une solution, mais celle-ci se révèle peu performante. Pour avoir un bon résultat, le calcul de cette probabilité fait nécessairement intervenir la loi de T . L'intervalle aléatoire $]T - \varepsilon, T + \varepsilon[$ est appelé **intervalle de confiance** pour τ de niveau $1 - \alpha$.

L'intervalle réel $]t - \varepsilon, t + \varepsilon[$ est « l'observation de l'intervalle de confiance » ou la « **fourchette** ». On ne sait pas avec certitude si τ est dedans.

b) Estimation d'une proportion p

Soit une population dans laquelle une modalité A d'un caractère qualitatif est en proportion p . Dans un prélèvement au hasard d'un élément de cette population, la probabilité qu'il présente la modalité A est donc p , valeur que l'on désire estimer.

Exemple : Un sondage effectué auprès de 150 personnes choisies de façon aléatoire dans une circonscription donne 45 suffrages au candidat A . Quelle est la proportion p des partisans de A dans la population ?

La valeur observée de la fréquence de A dans l'échantillon est ici : $f = 0,3$.

On choisit f comme estimation ponctuelle de la proportion p d'électeurs favorables au candidat A . Soit X_0 la variable de BERNOULLI qui prend la valeur 1 pour un électeur de A et 0 sinon. La variable $\sum X_i = n\bar{X}$ est égale au nombre des partisans de A dans l'échantillon. C'est une variable binomiale $B(n; p)$.

On sait que $E\left(\sum X_i\right) = np$, $\text{Var}\left(\sum X_i\right) = np(1 - p)$, d'où $E(\bar{X}) = p$ et $\text{Var}(\bar{X}) = \frac{p(1 - p)}{n}$.

$\bar{X}$ est donc un estimateur sans biais et convergent de p .

On cherche un intervalle de confiance de p de niveau $1 - \alpha = 0,95$. Pour avoir un bon résultat, il nous faut connaître la loi de $\bar{X}$, moyenne des X_i . Or la loi binomiale se prête mal aux calculs. On démontre (théorème de MOIVRE-LAPLACE) que la loi de la variable

réduite $U = \frac{\bar{X} - p}{\sqrt{\frac{\bar{X}(1 - \bar{X})}{n}}}$ peut être approchée par la loi normale^(f) $N(0; 1)$ pour n assez grand ($n \geq 50$).

La condition que doit vérifier l'intervalle de confiance de niveau $1 - \alpha$ est : $\mathbb{P}(\bar{X} - \varepsilon < p < \bar{X} + \varepsilon) = 1 - \alpha$, elle s'écrit

$$\mathbb{P}\left(\frac{-\varepsilon}{\sqrt{\frac{\bar{X}(1 - \bar{X})}{n}}} < U < \frac{\varepsilon}{\sqrt{\frac{\bar{X}(1 - \bar{X})}{n}}}\right) = 1 - \alpha.$$

α étant donné, pour chaque valeur prise par $\bar{X}$, la table de la loi normale centrée réduite donne une valeur u pour $\frac{\varepsilon}{\sqrt{\frac{\bar{X}(1 - \bar{X})}{n}}}$ qui vérifie la condition de confiance $\mathbb{P}(-u < U < u) = 1 - \alpha$. On obtient donc l'intervalle de confiance pour p :

$$\left[\bar{X} - u\sqrt{\frac{\bar{X}(1 - \bar{X})}{n}} ; \bar{X} + u\sqrt{\frac{\bar{X}(1 - \bar{X})}{n}} \right]$$

Avec l'observation f de $\bar{X}$, la demi-longueur de l'intervalle de confiance est alors $\varepsilon = u\sqrt{\frac{f(1 - f)}{n}}$ et la fourchette obtenue pour estimer la proportion p au niveau de confiance $1 - \alpha$ est donc :

$$\left[f - u\sqrt{\frac{f(1 - f)}{n}} ; f + u\sqrt{\frac{f(1 - f)}{n}} \right]$$

Dans l'exemple, on a la valeur observée $f = 0,3$ pour $\bar{X}$ et une taille $n = 150$ pour l'échantillon. Avec $\alpha = 0,05$, on obtient $u = 1,96$, ce qui donne la fourchette $]0,226 ; 0,374[$ pour estimer p au niveau de confiance 0,95.

Remarquons qu'en pourcentage, p est donnée à 7,4 % près avec 5 chances sur 100 de se tromper, ce qui n'est pas fameux pour un sondage. L'échantillon est trop petit pour estimer assez précisément cette proportion. En prenant $n = 1000$, on obtient la fourchette $]0,271 ; 0,329[$ susceptible d'encadrer p à environ 3 % près, ce qui est la performance habituelle des sondages médiatisés.

Remarquons aussi que la précision de cette estimation ne dépend pas de la taille de la population P mais qu'elle est inversement proportionnelle à la racine carrée de la taille n de l'échantillon sondé.

^(f)On utilise ici cette loi dont la connaissance n'est pas essentielle pour comprendre la suite.

Pour simplifier un peu grossièrement, on peut aussi majorer $\sqrt{f(1-f)}$ par 0,5 ce qui, en augmentant la valeur calculée pour ε , garantit un niveau de confiance supérieur à $1 - \alpha$. Cette majoration n'est pas trop brutale : pour $f = 0,3$, on a $\sqrt{f(1-f)} = 0,46$ et pour $f = 0,1$, on a $\sqrt{f(1-f)} = 0,3$.

Avec cette simplification, la condition $\mathbb{P}(|U| < u) \geq 1 - \alpha$ donne $\varepsilon < \frac{u}{2\sqrt{n}}$. Mais pour $\alpha = 0,05$, on a $u = 1,96$ d'où à peu près $\varepsilon = \frac{1}{\sqrt{n}}$, valeur proposée dans un thème du nouveau programme de statistique de seconde. On obtient, en fonction de différentes valeurs de n et de α , les demi-fourchettes ε suivantes en pourcentages :

$\alpha \backslash n$	500	800	1 000	2 000	10 000
0,1	3,7	2,9	2,6	1,8	0,8
0,05	4,4	3,5	3	2,2	①
0,01	5,7	4,5	4	2,9	1,3

Par exemple, il faudrait un sondage de 10 000 personnes pour donner la répartition d'un choix dans la population à 1 % près, avec encore 5 chances sur 100 de se tromper.