

ANTOINE BODIN

**Analyse implicative : modèles sous-jacents à l'analyse
implicative et outils complémentaires**

Publications de l'Institut de recherche mathématiques de Rennes, 1995-1996, fascicule 3
« Fascicule de didactique des mathématiques et de l'E.I.A.O. », , exp. n° 3, p. 1-23

http://www.numdam.org/item?id=PSMIR_1995-1996__3_A4_0

© Département de mathématiques et informatique, université de Rennes,
1995-1996, tous droits réservés.

L'accès aux archives de la série « Publications mathématiques et informatiques de Rennes » implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.

NUMDAM

Article numérisé dans le cadre du programme
Numérisation de documents anciens mathématiques
<http://www.numdam.org/>

Analyse implicative : modèles sous-jacents à l'analyse implicative et outils complémentaires

Antoine BODIN
IREM Besançon

Situation de base

Pour une raison ou pour une autre, un ensemble E d'individus est l'objet de notre attention. Ces individus peuvent ou non présenter les caractères que nous notons a et b . Ces caractères déterminent donc des variables binaires.

Nous noterons respectivement n_a et n_b les cardinaux des ensembles A et B d'individus *observés* qui présentent respectivement les caractères a et b .

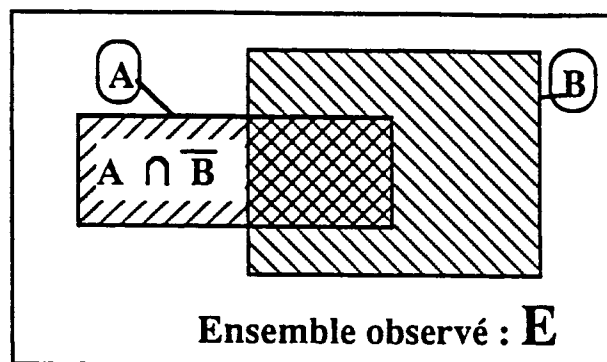


Figure 1

On peut être amené à se demander si une observation faite (la contingence) est compatible avec telle ou telle hypothèse qui pourrait être faite sur ce couple de caractères.

Par exemple :

- Les caractères a et b sont indépendants (ex. 1).
- Une variable aléatoire étant définie sur E , la moyenne arithmétique des éléments qui possèdent le caractère a est égale à la moyenne arithmétique des éléments qui possèdent b (ex. 2).
- Les individus qui possèdent le caractère a possèdent aussi le caractère b (ex. 3).
- Les individus qui possèdent le caractère a ont une certaine tendance à posséder aussi le caractère b (ex. 4).
- ...

Si l'on reste au niveau des observations, la réponse à ces questions est soit triviale (ex. 1, 2 et 3), soit susceptible d'interprétations multiples (ex 4).

Deux cas doivent alors être distingués :

1er cas : Observation exhaustive d'une population

Dans ce cas, E est la population mère, c'est à dire que nous ne cherchons pas, ni ne chercherons, à inférer *directement* les résultats que nous pourrions obtenir sur E , à une population plus vaste dont E serait issue.

Voir cependant la démarche du raisonnement plausible de Polya (Polya 1958), ou la méthode d'induction utilisée couramment aussi bien dans les sciences physiques que dans les sciences humaines.

2ème cas : E est un sous ensemble d'une population-mère

Dans ce cas, E est une réalisation d'un tirage effectué dans une population mère que nous noterons $\mathcal{E}$.

Bien sûr l'une des questions à se poser porte sur la possibilité de considérer E comme un échantillon de $\mathcal{E}$.

Les ensembles A et B observés, sont alors des parties, éventuellement des échantillons, des ensembles $\mathcal{A}$ et $\mathcal{B}$ des individus qui possèdent respectivement les caractères a et b (figure 2).

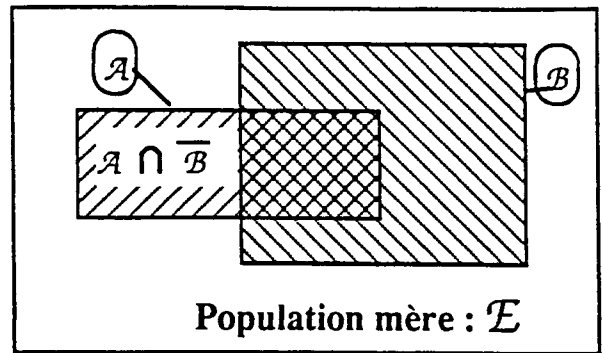


Figure 2

Questions fondamentales relatives à la situation de base.

Questions du premier type (de nature tendancielle) :

Est-il possible de considérer que la possession du caractère a tend à "impliquer" la possession du caractère b ? (Avec un sens du mot impliquer qui, pour l'instant, peut rester vague).

Questions du second type (de nature inclusives) :

Est-il possible de considérer que la possession du caractère a assure la possession du caractère b avec un certain degré de confiance ?

Nous allons voir que la théorie implicative telle qu'elle a été développée par Régis Gras, et par ses élèves, constitue une réponse efficace aux questions du premier type mais qu'elle ne répond pas complètement à celles du second.

Nous verrons aussi que c'est souvent des questions du second type qui, dans un premier temps, intéressent le chercheur. Cela explique sans doute pourquoi dans les travaux sur l'analyse implicative, la référence à l'inclusion revient tellement souvent (y compris dans la thèse initiale de R. Gras).

Nous montrerons plus loin que les deux types de questions doivent en fait être considérés comme strictement distincts et que les démarches qu'il est possible de développer pour y répondre seront tout aussi distinctes, bien qu'elles puissent être complémentaires.

La grande variété des outils statistiques mis à la disposition des chercheurs suppose de ceux-ci une connaissance précise des conditions de fonctionnement et de validité des modèles utilisés. Le présent chapitre a essentiellement pour but d'apporter quelques éclairages sur le modèle implicatif et le situant par rapport à des modèles alternatifs envisageables, et cela de façon à en préciser tout à la fois les forces et les limites.

Pour cela, nous allons commencer par expliciter les hypothèses sous-jacentes aux divers modèles qu'il est possible de construire pour répondre aux questions posées ci-dessus (questions fondamentales relatives à la situation de base).

Notations utilisées

Dans tous les cas, nous noterons :

n_a et n_b les cardinaux respectifs des ensembles A et B d'individus *observés* qui présentent respectivement les caractères a et b,

$p(a)$ et $p(b)$ les probabilités des caractères a et b *dans la population mère*,

$n = \text{card}(E)$,

$K_0 = \text{card}(A \cap \overline{B})$,

$p_1, p_2, p_3, \dots$ les différentes fonctions de probabilité que nous introduirons.

Θ désignant un variable aléatoire, nous noterons $E(\Theta)$ l'espérance mathématique de Θ .

Approches relatives aux questions de type tendanciel

I - Cas où $E = \mathcal{E}$

L'observation est alors exhaustive, c'est à dire que l'on connaît les cardinaux n_a , n_b , et $n_{a \wedge b}$.

I - 1 : Modèle de l'indépendance statistique

Dans ce cas, l'indépendance des caractères a et b est caractérisée

par la relation : $\frac{n_{a \wedge b}}{n_a} = \frac{n_b}{n}$.

Dans notre univers, cette relation *est* ou *n'est pas* vérifiée.

Si elle n'est pas vérifiée, on peut par exemple s'interroger sur la signification (ou sur l'importance) de l'écart $\frac{n_{a \wedge b}}{n_a} - \frac{n_b}{n}$.

Une façon d'aborder la question pourrait conduire à construire un indice i défini de façon telle que (cf. Figure 3) :

- si $\mathcal{A} \subset \mathcal{B}$, $i(a; b) = 1$ (implication logique),
- si a et b sont indépendants (c'est à dire si $p(b/a) = p(b)$), $i(a; b) = 1$
- i soit une fonction affine de la variable $x = p(b/a)$.

Pour cela on peut toujours supposer b fixé et a variable.

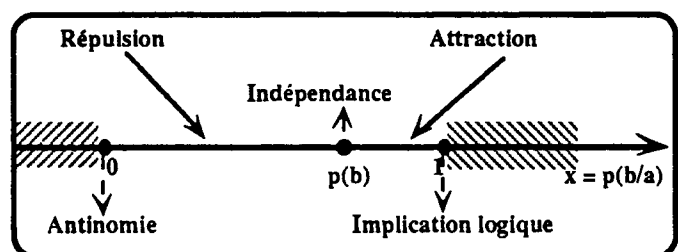


Figure 3

La linéarisation conduit à :

$$i(a ; b) = \frac{p(a \cap b) - p(a)p(b)}{p(a)(1 - p(b))} .$$

Une transformation élémentaire de cette relation conduit à :

$$i(a ; b) = 1 - \frac{p(a \cap \bar{b})}{p(a)p(\bar{b})} .$$

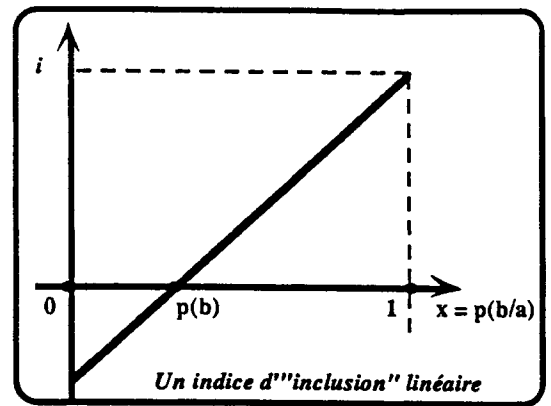


Figure 4

L'indice trouvé n'est autre que l'indice de Jane Loevinger, déjà signalé (sans justification) dans la thèse de R. Gras (Gras 1977).

Dans le cadre de la mise en perspective annoncée, il m'a semblé intéressant de rappeler ici cet indice en montrant sur quelle idée il pouvait être construit. Il correspond en effet à un modèle possible, parmi divers modèles envisageables. Cela dit, R. Gras a montré depuis longtemps l'insuffisance de cet indice dû en particulier au fait qui n'a aucune sensibilité à la taille de E.

I - 2 : Modèle de l'hypothèse d'absence de lien

On peut considérer que la contingence, ici, consiste en l'observation de ces n_a et n_b , individus, possédant respectivement les caractères a et b, plutôt qu'en n_a et n_b autres individus (dans ce même ensemble E de taille n).

Cela correspond *exactement* à l'hypothèse initiale (et fondatrice !) de R. Gras (Gras 1977), considérant, à la suite de I.C. Lerman, que l'on pouvait penser à deux ensembles X et Y de tailles connues n_a et n_b , qui pouvaient se mouvoir indépendamment dans E de taille connue n.

Posons $K_0 = \text{card}(A \cap \bar{B})$, et introduisons la variable aléatoire K définie par : $K = \text{card}(X \cap \bar{Y})$.

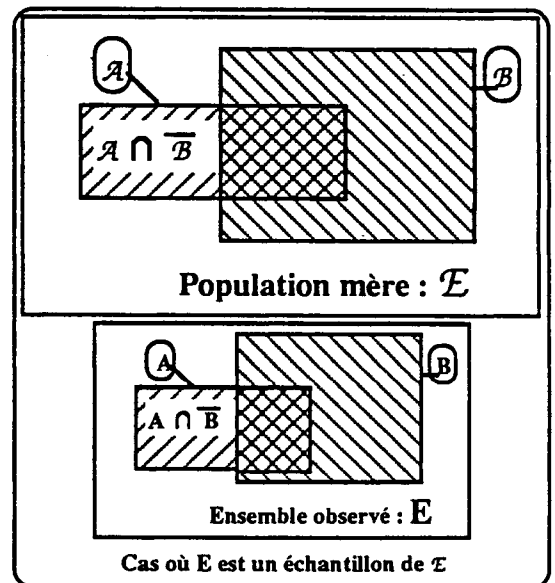


Figure 5

L'ensemble des événements élémentaires est l'ensemble Ω formé de l'ensemble des couples d'ensembles de E de cardinaux respectifs n_a et n_b .

$$\Omega = \{ (X, Y) ; X \subset E, Y \subset E, \text{card } X = n_a, \text{Card } Y = n_b \}$$

On a aussi :

$$X \in \mathcal{X} / \mathcal{X} = \{ T \subset E / \text{Card}(T) = n_a \}$$

$$Y \in \mathcal{Y} / \mathcal{Y} = \{T \subset E / \text{Card}(T) = n_b\}$$

Posons alors comme hypothèse nulle (H_0) l'équiprobabilité des éléments de Ω .

Sous H_0 :

le nombre de réalisations possibles

(équiprobables) est : $\binom{n_a}{n} \binom{n_b}{n}$,

le nombre de réalisation pour lesquelles

$K = K_0$ est : $\binom{n_b}{n} \binom{n_a - K_0}{n_b - K_0} \binom{K_0}{n - n_b}$,

et l'on a :

$$p_1(K = K_0) = \frac{\binom{n_b}{n} \binom{n_a - K_0}{n_b - K_0} \binom{K_0}{n - n_b}}{\binom{n_a}{n} \binom{n_b}{n}},$$

soit, en simplifiant par $\binom{n_b}{n}$:

$$p_1(K = K_0) = \frac{\binom{n_a - K_0}{n_b - K_0} \binom{K_0}{n - n_b}}{\binom{n_a}{n}}$$

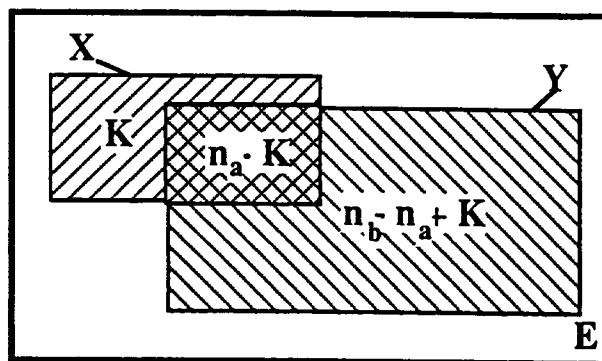


Figure 6

Ce dernier résultat montre que, sous H_0 , K suit la loi hypergéométrique de paramètres (n, n_a, p) ; p étant la proportion d'éléments vérifiant $\bar{B}$ dans E (proportion observée).

Le calcul précédent montre qu'il revient au même de considérer X et Y variables ou de laisser Y fixe et de laisser X varier.

Construction du test d'hypothèse

L'hypothèse nulle étant posée comme précisé ci-dessus (équiprobabilité des éléments de Ω , ou, ce qui revient au même, absence de lien),

$$\text{sous } H_0 : E(K) = E(\text{card}(X \cap \bar{Y})) = n_a \times \frac{n_b}{n}.$$

Puisque nous nous intéressons à la tendance des éléments qui possèdent le caractère a à posséder aussi le caractère b nous poserons comme hypothèse alternative, l'hypothèse suivante :

$$H_1 : E(\text{card}(X \cap \bar{Y})) < n_a \times \frac{n_b}{n}$$

C'est à dire qu'en moyenne nous nous attendons à trouver moins d'éléments dans $A \cap \bar{B}$ que ce que l'absence de lien laisserait prévoir.

Il revient au même de dire que l'on s'attend à trouver plus d'éléments dans $A \cap B$ que ce qui est observé.

D'où le test :

Après avoir choisi un niveau de risque de première espèce $1 - \alpha$ (c'est à dire un seuil de confiance α),

si $p_1(K \geq K_0) < \alpha$, on ne rejette pas l'hypothèse H_0 , au seuil de confiance α ,

si $p_1(K \geq K_0) > \alpha$, on ne rejette pas H_1 , au seuil de confiance α (en fait on l'adopte, au moins temporairement),

dans ce cas, on dit que **a** implique **b** au seuil de confiance α (dans le modèle de l'absence de lien).

En posant : $\varphi_1 = (1 - p_1(K \leq K_0))$,

on trouve un candidat possible pour être un indice d'implication. Cet indice a été envisagé un moment.

L'avantage de cette approche est évidemment qu'elle reste centrée sur le groupe observé (considéré comme population) et dont il n'est pas besoin de se demander s'il constitue ou non un échantillon d'une quelconque population mère. L'inconvénient corrélatif réside bien sûr dans le risque d'extrapolations abusives.

Insistons sur le fait que, dans cette approche, l'hypothèse nulle (H_0) concerne l'équiprobabilité des éléments de Ω . Si l'on tient à parler d'indépendance dans ce cas, il faudrait parler de l'indépendance des tirages des ensembles X dans $\mathcal{X}$ et Y dans $\mathcal{Y}$.

II - Cas où E peut être considéré comme étant un échantillon de $\mathcal{E}$, (E de taille fixée)

II - 1 : Cas où la population mère est de taille N finie (modèle hypergéométrique)

Ici, N , n sont supposés connus tandis que n_a , n_b , ainsi que $n_{a \wedge b}$ sont des valeurs observées.

La seule hypothèse concerne l'indépendance de **a** et **b** dans $\mathcal{E}$.

Cette hypothèse qui constituera l'hypothèse nulle, s'écrit :

$$H_0 : p(b/a) = p(b)$$

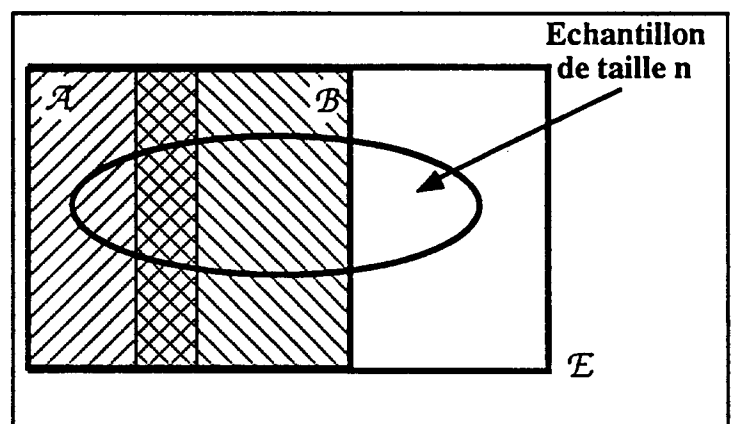


Figure 7

Nous prendrons alors comme hypothèse alternative :

$$H_1: p(b/a) > p(b)$$

L'expérience consiste alors simplement à extraire de $\mathcal{E}$ un échantillon E de taille n .

Dans ces conditions, K suit alors la loi **hypergéométrique** de paramètres (N, n, p_2) , p_2 étant la proportion d'éléments vérifiant a et non b dans $\mathcal{E}$.

$$\text{D'où : } p_2 = \frac{\text{card}(\mathcal{A} \cap \overline{\mathcal{B}})}{N}.$$

L'hypothèse H_0 se traduit ici par : $p_2 = \frac{N_a}{N} \times \frac{N_{\overline{b}}}{N}$.

Contrairement à ce qui se passait dans le modèle issu de l'hypothèse d'absence de lien (modèle I-2), ici, on ne connaît pas p .

La démarche du maximum de vraisemblance nous amène à estimer p_2 par :

$$\frac{n_a}{n} \times \frac{n_{\overline{b}}}{n}.$$

$$\text{Sous } H_0 : E(K) \approx n \times \frac{n_a}{n} \times \frac{n_{\overline{b}}}{n}.$$

$$\text{Ou encore, } E(K) \approx \frac{n_a \times n_{\overline{b}}}{n}.$$

Le test d'hypothèse s'effectue de la même façon que précédemment :

si $p_2(K \geq K_0) < \alpha$, on ne rejette pas l'hypothèse H_0 , au seuil de confiance α ,

si $p_2(K \geq K_0) > \alpha$, on ne rejette pas H_1 , au seuil de confiance α (en fait on l'adopte, au moins temporairement),

Dans ce cas, on dit que a implique b au seuil de confiance α (dans le modèle hypergéométrique).

et l'on obtient un autre candidat, φ_2 , pour l'indice d'implication :

$$\varphi_2 = (1 - p_2(K \leq K_0))$$

Bien que les tests d'hypothèses construits en I-2 et en II-1 semblent identiques (à partir des mêmes observations, ils conduisent, de fait, aux mêmes décisions), les modèles qui les sous-tendent sont différents. On se gardera d'oublier, par exemple, que dans l'un des cas nous utilisons un paramètre estimé, ce qui n'est pas le cas dans l'autre.

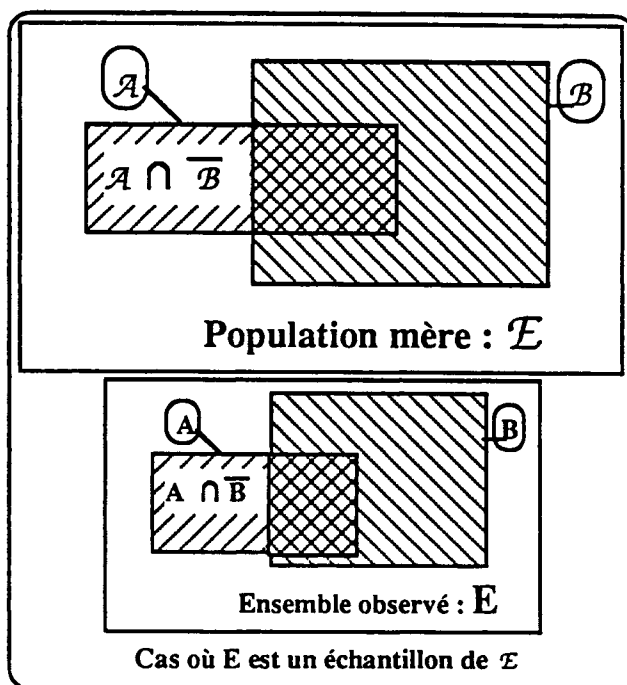


Figure 8

II - 2 : Cas où la population mère est infinie (modèle binomial)

Posons encore comme hypothèse nulle (H_0) l'indépendance des caractères a et b , soit :

$$H_0 : p(a \wedge b) = p(a).p(b) ,$$

et, puisque nous nous intéressons toujours à la tendance de a à être b , posons comme hypothèse alternative :

$$H_1 : p(a \wedge b) > p(a).p(b),$$

ou, ce qui est équivalent : $H_1 : p(a \wedge \bar{b}) \leq p(a).p(\bar{b})$.

Soit toujours K_0 le cardinal de $A \cap \bar{B}$ (observé), et K la variable aléatoire correspondante, pour un tirage de taille n , au nombre d'éléments vérifiant a et $\bar{b}$.

Sous H_0 , K suit la loi binomiale de paramètres n et $p(a).p(\bar{b})$.

Notons $p_3 = p(a).p(\bar{b})$.

Comme dans le cas précédent (II-1) on ne peut qu'estimer les probabilités $p(a)$ et $p(\bar{b})$.

Nous prendrons encore $\frac{n_a}{n}$ et $\frac{n_{\bar{b}}}{n}$ comme estimations de $p(a)$ et $p(\bar{b})$.

Le test d'hypothèse au seuil α , s'effectue toujours de la même façon :

si $p_3(K \geq K_0) < \alpha$, on ne rejette pas l'hypothèse H_0 , au seuil de confiance α .,

si $p_3(K \geq K_0) > \alpha$, on accepte l'implication de a vers b , au seuil de confiance α .,

Dans ce cas, on dit que a implique b au seuil de confiance α (*dans le modèle binomial II-2*).

En posant : $\phi_3 = (1 - p_3(K \leq K_0))$

on trouve à nouveau l'indice construit par Régis Gras et utilisé dans sa thèse (Gras, 1979).

Comme le calcul exact de cet indice est difficile pour les valeurs de n qui nous intéressent (loi hypergéométrique !), l'indice était alors calculé en utilisant l'approximation normale.

Il convient toutefois de remarquer que l'approche faite à ce moment par Régis Gras, bien que conduisant au même indice, correspondait à un modèle différent de celui présenté ici.

III - Cas où E peut être considéré comme étant un échantillon de $\mathcal{E}$, de taille indéterminée.

III - 1 : Modèle de Poisson

La population mère étant toujours infinie, supposons que l'on ait tiré un échantillon dont la taille elle-même est aléatoire.

Notons m cette taille.

Supposons que nous nous intéressions toujours au nombre K d'éléments qui vérifient les caractères a et non b . Soit p_4 la probabilité correspondante.

Avec les notations utilisées dans I et II, nous avons :

$$p_4(K = s) = \sum_{N=0}^{\infty} p(m = N) \times p[(K = s) / (m = N)],$$

ce qui s'écrit encore :

$$p_4(K = s) = \sum_{n \geq s} p(n = m) \times (p(K = s) / (n = m)).$$

Supposons que la variable m suive une loi de Poisson. La valeur observée de m étant n , nous prendrons n comme estimateur du paramètre μ de cette loi.

Posons encore comme hypothèse nulle l'indépendance de a et de b .

Le modèle ainsi introduit est exactement celui utilisé par Annie Larher dans sa thèse (Larher, 1991). Dans ce cas, elle démontre (page 5) que la variable K suit la loi de Poisson de paramètre $\lambda = n.p(a).p(\bar{b})$.

Pour mémoire, voici cette démonstration.

En posant : $\pi = p(a \wedge \bar{b})$, nous avons :

$$p(n = m) = \frac{n^m}{m!} e^{-n} \quad \text{et} \quad p_4(K = s) = \sum_{m \geq s} p(n = m) \times (p(K = s) / (n = m))$$

$$\text{et donc : } p_4(K = s) = \sum_{m \geq s} \frac{n^m}{m!} e^{-n} \times \left(\int_m^s \pi^s (1 - \pi)^{m-s} \right)$$

Soit :

$$p_4(K = s) = \frac{(n\pi)^s}{s} e^{-n\pi} \sum_{m \geq s} \left(\frac{(n(1 - \pi))^{(m-s)}}{(m-s)!} e^{-n(1-\pi)} \right)$$

et :

$$\sum_{m \geq s} \left(\frac{(n(1 - \pi))^{(m-s)}}{(m-s)!} e^{-n(1-\pi)} \right) = \sum_{m \geq 0} \left(\frac{(n(1 - \pi))^m}{m!} e^{-n(1-\pi)} \right)$$

Le second membre est la somme des termes d'une distribution de Poisson de paramètre $n(1 - \pi)$.

$$\text{Donc : } \sum_{m \geq s} \left(\frac{(n(1 - \pi))^{(m-s)}}{(m-s)!} e^{-n(1-\pi)} \right) = 1$$

$$\text{Et : } p_4(K = s) = \frac{(n\pi)^s}{s} e^{-n\pi}$$

K suit ainsi la loi de Poisson de paramètre $\lambda = n \pi$.

L'hypothèse d'indépendance de a et b se traduit encore par : $p(a \wedge \bar{b}) = p(a).p(\bar{b})$

D'où : $\lambda = n.p(a).p(\bar{b})$

Comme dans les parties I et II ci-dessus, $p(a)$ et $p(\bar{b})$ sont estimées, à partir des valeurs observées, par $\frac{n_a}{n}$ et $\frac{n_{\bar{b}}}{n}$.

Finalement, cela conduit à estimer π par $\frac{n_a \cdot n_{\bar{b}}}{n}$.

Ici, l'hypothèse nulle doit être ainsi formulée :

H_0 : N suit une loi de Poisson *et* a et de b sont indépendants.

Le test d'hypothèse au seuil de confiance α s'effectue encre de la même façon que dans les cas précédents :

si $p_4(K \geq K_0) \leq \alpha$, on ne rejette pas l'hypothèse H_0 , au seuil de confiance α .,

si $p_4(K \geq K_0) > \alpha$, on accepte l'implication de a vers b , au seuil de confiance α .,

En posant :

$$\varphi_4 = (1 - p(K \leq K_0)),$$

on trouve l'indice d'implication φ_4 introduit dans la thèse d'A. Larher et utilisé depuis (en fait, c'est une approximation par la loi normale qui est utilisée).

III - 2 : Autre présentation possible : modèle du flux poissonnien

La population mère étant toujours infinie, supposons que l'on tire, séquentiellement, des éléments de l'ensemble $\mathcal{E}$. Supposons aussi que l'on note les instants de tirage des éléments. Remarquons que cela correspond à une situation plausible pour un correcteur de copies ou pour tout autre observateur séquentiel.

Supposons que ce soit encore l'événement " a et non b " qui retienne spécialement notre attention. Notons Y l'événement "*un élément est observé*", et X l'événement "*un élément est observé et il vérifie a et non b* ".

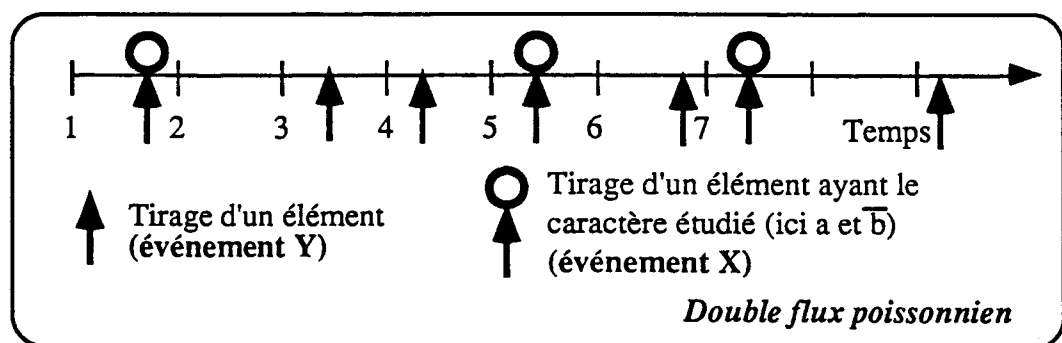


Figure 9

Nous avons alors deux variables aléatoires :

Y_t : Le nombre d'événements Y observés pendant la période $[0 ; t]$,

X_t : Le nombre d'événements X observés pendant la période $[0 ; t]$.

Faisons alors les hypothèses suivantes :

- 1 - Le processus d'attribution du caractère X (resp. Y) est sans mémoire (attention aux effets d'attente en évaluation).
- 2 - La loi du nombre d'événement X (resp. Y) relevés dans un intervalle $[t_0, t_0 + t]$ ne dépend que de t .
- 3 - Deux événement X (resp. Y) ne peuvent se produire simultanément.

Soit p (resp. c) la cadence de l'événement X (resp. Y) : p (resp. c) est l'espérance mathématique du nombre de tels événements relevés dans un intervalle de temps d'une unité.

$E(X_1) = p$ (avec nos notations, X_1 est la variable aléatoire égale au nombre d'évènements X relevés pendant une unité de temps).

$E(Y_1) = c$ (avec nos notations, Y_1 est la variable aléatoire égale au nombre d'évènements Y relevés pendant une unité de temps).

Dans ces conditions, on démontre que X_t (resp. Y_t) suit une loi de Poisson de paramètre $\lambda = pt$ (resp. ct).

Le lecteur trouvera une démonstration dans SAPORTA (page 35) et une autre dans ENGEL (volume 2 page 237)

D'où : $E(X_t) = pt$; $E(Y_t) = ct$

Soit T une durée donnée (période de temps), qui sera justement notre durée, supposée fixe, d'observation.

Supposons, ce qui ne change pas la généralité, que la cadence des observations soit $c = 1$. Cela signifie que l'on fait en moyenne une observation par unité de temps.

Nous avons finalement : $\lambda = E(X_T) = pT$ et $T \approx n$.

Posons toujours comme hypothèse nulle H_0 l'indépendance de a et de b .

Comme précédemment, sous H_0 on peut estimer p par $\frac{n_a}{n} \times \frac{n_b}{n}$.

Ce qui conduit encore à estimer λ par $\frac{n_a \cdot n_b}{n}$.

Le test d'hypothèse se construit alors *exactement* comme dans le cas précédent (III-1).

On obtient le même indice d'implication (ϕ_4) que précédemment, mais avec une hypothèse sous-jacente peut-être plus "réaliste".

Insistons sur le fait que, bien que le test d'hypothèse soit *apparemment* le même que dans le cas III.1, et que l'indice d'implication retenu soit *exactement* le même, les modèles sous-jacents sont différents.

Loi de Poisson ou loi binomiale ?

La façon dont la loi de Poisson est habituellement sollicitée dans la théorie implicative peut être soumise à discussion.

Le modèle utilisé suppose en effet le tirage au hasard d'un échantillon de taille elle-même aléatoire. Cette hypothèse n'est pas sans intérêt dans les situations où la taille de la sous-population étudiée n'a pas été l'objet d'un choix préalable : par exemple lorsque l'on étudie les résultats d'une classe particulière, le nombre d'élèves concernés qui par exemple peut être 28, aurait tout aussi bien être 29, 27, ...

Il n'en va pas de même lorsque le choix de la taille de l'échantillon est fait a priori. Par exemple dans les études à grande échelle, on dispose d'un nombre très grand de cas et, pour les analyses, on est amené à échantillonner. La taille de l'échantillon est dans ce cas fixé a priori et on ne voit plus pourquoi il serait utile de faire l'hypothèse qui conduit à (modèle III 1):

$$p(K = K_0) = \sum_{N=0}^{\infty} p(n = N) \times (p(K = K_0)/(n = N))$$

Rappelons que dans ce cas, la loi de K est une loi de Poisson de paramètre λ , donc d'espérance et de variance toutes les deux égales à λ .

Nous avons vu que, dans ce cas, le paramètre λ pouvait être estimé par $\frac{n_a n_b}{n}$,

Dans le cas du modèle binomial (modèle II-2), K suit simplement la loi binomiale de paramètres n (taille de l'échantillon) et p . Ici, p peut être estimé par $\frac{n_a}{n} \times \frac{n_b}{n}$.

L'espérance de K est alors $E(K) = n.p$, et sa variance est $V(K) = n.p.(1 - p)$, valeurs qui peuvent être approchées, respectivement par

$$\frac{n_a n_b}{n} \quad \text{et} \quad \frac{n_a n_b}{n} \left(1 - \frac{n_a n_b}{n^2}\right)$$

Dans ces deux cas (I-1 à III-2) l'hypothèse nulle est celle de l'indépendance des caractères a et b . Dans les deux cas la moyenne de K est égale à λ , mais dans le cas binomial la variance est :

$$V(K) = \lambda (1 - p)$$

Cette variance est donc inférieure à celle calculée dans les conditions du modèle de Poisson.

Dans le cas où le modèle binomial serait pertinent, on remarque que $p(K \leq K_0)$ est surestimée lorsque l'on utilise la loi de Poisson plutôt que la loi binomiale. Dans les mêmes conditions, l'indice d'implication retenu est sous-estimé.

Cette surestimation allant en sens contraire de l'implication, elle ne constitue en fait qu'une assurance supplémentaire de la qualité des implications statistiques retenues...

On trouvera en annexe 1 une table de comparaison des indices correspondant aux modèles "absence de lien", "binomial" et "Poisson", ces indices étant calculés de façon exacte lorsque cela a été possible, et de façon approchée dans les autres cas.

Comme prévu, le modèle de Poisson est toujours le plus sévère pour la décision d'implication à quelque seuil que ce soit. On voit cependant le peu de différence qu'il y a en ce qui concerne les décisions prises sous l'un ou l'autre des modèles évoqués.

Pour une taille N petite (inférieure à 100), il pourrait y avoir intérêt à utiliser le modèle binomial qui évite de négliger trop d'implications acceptables (et que le modèle de Poisson refuserait).

Si l'on pense nécessaire de prendre davantage de précautions, c'est bien le modèle de Poisson qui garde notre préférence.

Considérations pratiques

L'objet de ce chapitre était une mise à plat des démarches inductives des modèles utilisés. Ceci étant fait, la légitimité de l'usage de l'indice d'implication et de la règle de décision, issus du modèle III-1 (ou III-2), et qui se sont imposés depuis le travail d'A. Larher s'en trouve renforcé et, sauf cas particulier évoqué ci-dessus, il n'y a que des avantages à continuer à les utiliser.

Conformément à cet usage nous désignerons maintenant par $\varphi(a, \bar{b})$ l'indice d'implication de a vers b et par p la probabilité correspondante.

Dans ce cas la distribution de probabilités correspondant converge avec n vers la loi normale $\mathcal{N}(\mu; \sigma)$, dont les paramètres peuvent être estimés à partir des observations :

$$\begin{aligned}\mu &\text{ peut être estimé par } \frac{n_a \cdot n_{\bar{b}}}{n} \\ \sigma &\text{ peut être estimé par } \sqrt{\frac{n_a \cdot n_{\bar{b}}}{n}}\end{aligned}$$

Dans le logiciel sous DOS développé par S. Ag Almouloud, l'approximation normale était systématiquement utilisée. Dans le logiciel sous Window dérivé du précédent (Bodin A & Couturier R, 1996), à temps de calcul égal, le calcul exact est toujours privilégié.

Ainsi, dans la version actuelle (8/96), le calcul exact est effectué pour $n_a \wedge \bar{b} < 48$.

De plus le calcul est toujours fait de façon exacte lorsque $\frac{n_a \cdot n_{\bar{b}}}{n} \leq 3$.

Modèles mathématiques et modèles fonctionnels

Les divers modèles mis à jour ci-dessus correspondent à des façons différentes de penser la situation et de construire des analogies productives de modèles dans lesquels le calcul sera possible.

Mais sans doute convient-il de s'interroger ici sur la notion de modèle.

En particulier, lorsque deux analogies différentes, et a priori non réductibles l'une à l'autre, débouchent sur le même modèle mathématique (un nombre, une fonction, une distribution, un ensemble,...), peut-on considérer qu'il s'agit du même modèle ? Ou plutôt, convient-il de gommer la différence au niveau des analogies initiales ?

Nous proposons de distinguer modèle mathématique et modèle fonctionnel.

Par exemple les analogies relatives respectivement à l'hypothèse d'absence de lien (modèle I-2) et à l'indépendance des caractères a et b définis sur un ensemble fini (modèle II-1)

conduisent au même modèle mathématique, tout en étant susceptible d'être pertinentes dans des situations très différentes. D'un point de vue fonctionnel, ces modèles sont donc différents.

Il en est de même des analogies conduisant au modèle de la distribution de Poisson chez A. Larher, et de celle du modèle du double flux poissonnien proposé ici.

Les modèles fonctionnels doivent donc être considérés comme distincts même dans les cas où les modèles mathématiques associés sont identiques.

Toutefois, une fois démontré l'équivalence des deux démarches (des deux analogies), on pourra dire que la seconde démarche fournit un modèle fonctionnel équivalent au premier (éventuellement plus simple).

Par exemple, l'analogie consistant à faire varier deux ensembles X et Y de tailles fixes respectives n_a et n_b est équivalente à l'analogie qui consiste à laisser Y fixe (de taille n_b) et à faire varier X (de taille n_a).

Implication et inclusion

Soit $\varphi(a, \bar{b})$ l'indice d'implication de a vers b .

On sait que φ est une mesure de "l'étonnement" d'obtenir "si peu" d'éléments appartenant à $a \cap \bar{b}$.

On a vu plus haut que le test maintenant classique de l'implication statistique au seuil α , peut se ramener à un test unilatéral d'indépendance.

Dans le cas où les effectifs sont petits (et même dans d'autres cas particuliers), même l'inclusion stricte de

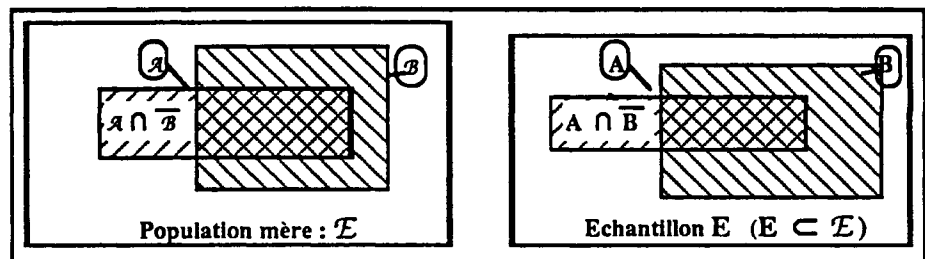


Figure 10

A dans B peut être due aux fluctuations d'échantillonnage et ne peut être considérée comme une manifestation d'un lien entre a et b . Dans ce cas, l'intérêt de l'indice d'implication est évident.

Cependant, lorsque n_a et n_b sont très grands, l'implication statistique, à quelque seuil que ce soit, se confond quasiment avec l'observation directe d'une proportion de $A \cap \bar{B}$ inférieure à son espérance mathématique dans le cas d'indépendance (elle se confond avec elle... asymptotiquement).

Dire que a implique b à un certain seuil α revient à dire que la proportion d'éléments qui vérifient simultanément a et non b est *moindre* que celle que l'on s'attendait à trouver sous l'hypothèse d'indépendance déjà évoquée.

Dans ces conditions on est loin d'une mesure de la quasi-inclusion de A dans B qui a été parfois évoquée. Il vaudrait mieux parler de tendance ou d'attraction de a par b , ou encore, de propension de a vers b (J.B. Lagrange). Il n'y a toutefois aucun d'inconvénient à continuer à

parler d'implication, mais à la condition stricte que toute confusion avec l'inclusion ou la quasi inclusion soit soigneusement évitée.

Ces remarques étant faites, on ne s'étonnera plus d'avoir simultanément

$$a \overset{\alpha}{\Rightarrow} b \quad \text{et} \quad \text{card}(A \cap \overline{B}) > \text{card}(A \cap B),$$

$$\text{et même : } \text{card}(\mathcal{A} \cap \overline{\mathcal{B}}) > \text{card}(\mathcal{A} \cap \mathcal{B})$$

Ce qui bien sûr s'accommode assez mal de l'idée d'inclusion de $\mathcal{A}$ dans $\mathcal{B}$.

Exemple

Supposons un ensemble E et deux parties A et B de E, telles que les cardinaux de E, A, B, $A \cap B, \dots$ soient dans les proportions du tableau ci-contre.

		a		
		1	0	
b	1	1	3	4
	0	2	14	16
		3	17	20

Pour préciser, voici deux réalisations possibles de cette situation, et d'autres cas intermédiaires sont données dans le tableau de valeurs numériques.

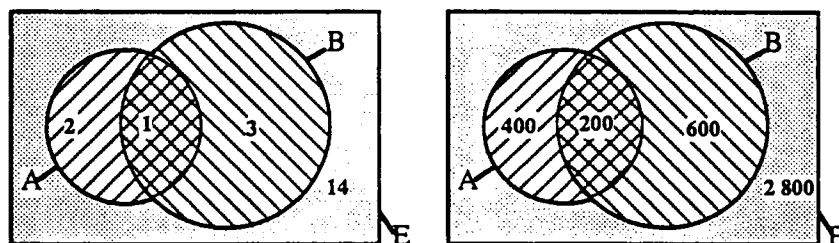


Figure 11

Dans chaque cas nous avons calculé la valeur du coefficient d'implication ϕ .

n	20	100	400	800	1600	3200	4000
n_a	3	15	60	120	240	480	600
n_b	4	20	80	160	320	640	800
$n_{a \text{ et } b}$	1	5	20	40	80	160	200
$n_{a \text{ et non } b}$	2	10	40	80	160	320	400
$\phi(a, \overline{b})$	0,602	0,718	0,876	0,949	0,99	0,999	0,9999

Figure 12

Dans tous les cas, les éléments qui ont le caractère a et qui n'ont pas le caractère b représentent les $\frac{2}{3}$ des éléments de A. Il est exact que compte tenu des tailles relatives de E, A et B on s'attendrait à ce que cette proportion soit supérieure à $\frac{2}{3}$. Sous réserve d'indépendance, elle devrait être de $\frac{20}{28}$ (environ 70%), ce qui explique les valeurs trouvées pour ϕ , dès que n est assez grand, et justifie, dans ces cas, l'énoncé de la conclusion sous l'une des formes :

"a est attiré par b",

"les éléments qui ont le caractère a ont tendance à avoir le caractère b"

ou mieux :

"les éléments qui ont le caractère a ont davantage tendance que les autres à avoir les caractère b", à condition d'explicitier le sens de cette phrase.

Dans la situation ci-dessus, si l'on devait faire un pari exclusif sur **b** sachant **a**, c'est bien sûr **non b** qui doit avoir notre préférence. Le fait de savoir, en plus, que $a \xrightarrow{\alpha} b$, n'est pas susceptible de modifier cette décision. Dans ce cas, il est d'ailleurs souvent plus intéressant d'introduire directement la variable **non b** et donc d'envisager l'implication $a \xrightarrow{\alpha} \text{non b}$.

Toutefois dans la situation précédente, le jeu équitable sachant **a** consisterait à **parier b à 1 contre 4**.

Si l'on sait de plus que $a \xrightarrow{\alpha} b$ alors le jeu devient favorable à qui **parie b à 1 contre 4** (toujours sachant a).

C'est comme cela qu'il convient de lire l'implication statistique. Mais il faut bien convenir que rien ne nous oblige à faire un tel pari, ou encore que l'observation de telle ou telle attirance peut ou non présenter un intérêt selon les situations. Il conviendra donc de garder ces remarques à l'esprit lors des interprétations.

Retour à la quasi-inclusion

Comme nous l'avons déjà signalé, ce qui intéresse bien souvent le chercheur c'est une mesure de l'écart à l'inclusion vraie de $\mathcal{A}$ dans $\mathcal{B}$.

De ce qui précède on peut sans doute conclure que l'indice d'implication n'est pas un bon candidat pour cette mesure.

En présence de deux caractères **a** et **b** tels que $n_a < n_b$, la question :

"Dans quelle mesure peut on compter sur le comportement b lorsque le comportement a se manifeste" garde tout son intérêt.

C'est pour tenter de répondre à cette question que j'ai été amené à proposer un test complémentaire.

Test d'hypothèse complémentaire

Dans bien des cas, et en particulier dans les recherches en didactique il serait intéressant de pouvoir répondre à la question :

τ étant un taux fixé à priori,

"peut-on estimer que parmi les individus qui, dans la population mère, ont le caractère a, le taux de ceux qui n'ont pas aussi le caractère b est inférieur ou égal à τ ".

Dans les travaux sur l'évaluation, avec $\tau = 5\%$ ou $\tau = 10\%$, on retrouve les taux classiquement associés aux erreurs d'inattention, et lapsus divers, qu'il est préférable de laisser de côté dans une première étude.

Avec τ de l'ordre de 30% on a quelque chance de se trouver en présence d'un comportement qui pourrait être associé à deux démarches différentes (pouvant correspondre à deux conceptions différentes), qui, en termes d'observables, ne se différencient pas au niveau de a mais qui se différencient au niveau de b.

Nous proposons alors de compléter l'analyse implicative par un test d'hypothèse simple :

On pose comme hypothèse nulle :

$H_0 : \tau \leq \tau_0$ (avec par exemple $\tau_0 = 05\%$, 01 % .. 10%,)

L'hypothèse alternative que nous retiendrons étant : $H_1 : \tau > \tau_0$

On considère alors l'observation faite comme une réalisation de l'expérience d'une épreuve binomiale de paramètres n et $\tau_0 \cdot p(a)$ (proportion maximum d'élément de $\mathcal{E}$ qui sous l'hypothèse nulle devraient vérifier a et non b).

De façon classique, nous approcherons cette loi binomiale par la loi normale de même moyenne et de même écart type.

K et K_0 étant définis comme précédemment,

pour $\tau = \tau_0$, et au seuil de confiance de $1 - \alpha$, nous avons :

$$K < n \cdot (\tau_0 \cdot p(a)) + \gamma_\alpha \sqrt{n \cdot (\tau_0 \cdot p(a)) \cdot (1 - \tau_0 \cdot p(a))}$$

γ_α étant la valeur fractale de la distribution normale relative à α , valeur définie par :

$$\Pr(u \leq \alpha) = \gamma_\alpha.$$

Soit : $\gamma_{0,99} \approx 2,32$; $\gamma_{0,95} \approx 1,65$;

Sous H_0 (ie pour $\tau \leq \tau_0$), on peut donc encore affirmer avec un confiance supérieure à $1 - \alpha$ que :

$$K < n \cdot (\tau_0 \cdot p(a)) + \gamma_\alpha \sqrt{n \cdot (\tau_0 \cdot p(a)) \cdot (1 - \tau_0 \cdot p(a))}$$

Comme précédemment, nous prendrons $\frac{n_a}{n}$ comme estimateur de $p(a)$.

Le test d'hypothèse proposé est le suivant :

On choisit τ_0 et α .

Pour la valeur de τ_0 choisie, et sous H_0 , on calcule l'intervalle d'acceptation de K au seuil de confiance $1 - \alpha$.

Soit $\mu = \mu(\tau_0, \alpha)$ la valeur supérieure de cet intervalle.

Si $K_0 \leq \mu$, on ne rejette pas H_0 , au seuil α (en fait on l'adopte, au moins temporairement).

Si $K_0 > \mu$, on rejette pas H_0 , et on adopte au seuil α et on adopte l'hypothèse alternative H_1 .

De même que nous avons noté $a \xrightarrow{\alpha} b$ pour exprimer l'implication statistique de a vers b au seuil α , nous noterons $\mathcal{A} \xrightarrow[1-\tau_0]{\alpha} \mathcal{B}$, pour exprimer que le test d'hypothèse défini dans cette section

n'a pas conduit à rejeter $H_0 : \tau \leq \tau_0$ au seuil de confiance α choisi.

$\mathcal{A}$ et $\mathcal{B}$ désignant toujours les ensembles associés aux variables a et b , nous dirons dans ce cas que $\mathcal{A}$ est inclus dans $\mathcal{B}$ à $(1 - \tau_0)\%$ (inclusion statistique).

EXEMPLES

Exemple 1 : étude ESIEE91

Dans l'étude ESIEE91 portant sur 1819 candidats à une École d'Ingénieurs, les résultats de l'analyse implicative appliquée au questionnaire QCM en 24 questions nous avaient surpris par le nombre relativement faible des implications relevées, que ce soit au seuil de 0,95 ou au seuil de 0,99.

Sur les 24 questions mettons momentanément de côté les deux questions suivantes :

- la question la mieux réussie (Q6 réussie par 72% des candidats), très impliquée et très incluyente,
- la question la moins bien réussie (Q12 réussie par 01% des candidats) et qui résiste assez bien au test d'inclusion mais peu impliquante.

Ces questions font l'objet d'une analyse séparée (voir étude ESIEE), disons simplement que la première est une question de logique, à contenu mathématique vide.

Cette réduction faite :

le nombre de couples (a, b) vérifiant $a \xrightarrow{.99} b$ n'est que 26,

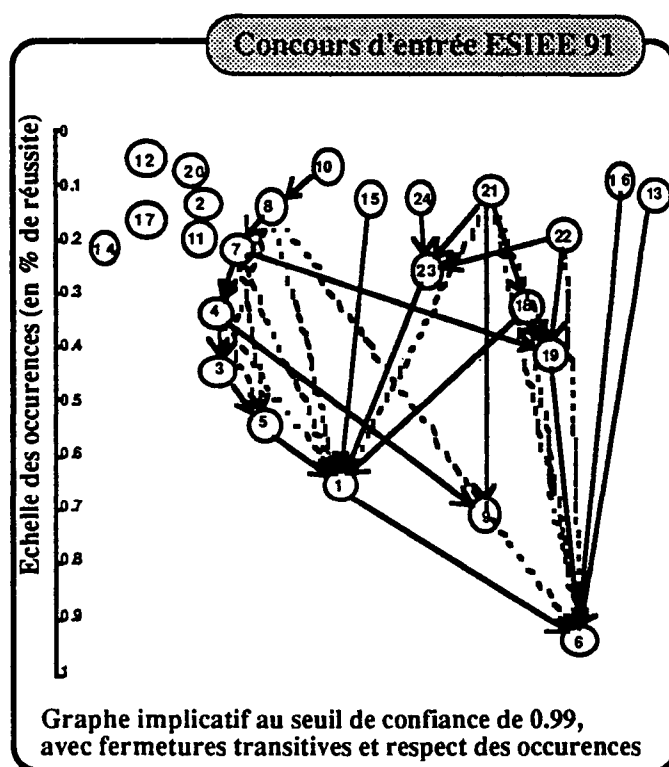


Figure 13.

et le nombre de couples (a, b) vérifiant $a \xrightarrow{.95} b$ n'est encore que 64, en comptant, bien sûr les arcs résultant de transitivités.

Ces nombres qui peuvent paraître importants nous avaient semblé faibles par rapport à ceux obtenus dans d'autres études a priori comparables que nous avons faites.

Au seuil de confiance 0,99 il apparaît une seule chaîne implicative de longueur supérieure à 3, il s'agit la chaîne :

10 - 8 - 7 - 4 - 3 - 5 - 1

Chaîne qui intègre deux nouveaux éléments si l'on adopte le seuil de confiance de 0,95, pour devenir :

10 - 8 - 21 - 7 - 4 - 3 - 5 - 1 - 9

De plus, toutes les sous-chaînes sont transitives au seuil de 0,95, sauf celles issues de 10, et cela pour des raisons de faible occurrence de cet item (05%).

On pouvait alors se demander si cette chaîne, a priori une bonne candidate pour constituer une échelle de repérage des compétence (Bodin, 1996), pouvait être considérée comme étant de type simplement tendancielle, ou s'il était possible de la considérer comme une pseudo-échelles de Guttman. Pour nous cela supposerait que cette chaîne soit aussi de type inclusif pour un choix de τ_0 au plus égal à 0,20 (valeur arbitraire qui peut être remise en question).

Le test d'hypothèse présenté dans cette partie a alors été utilisé pour différentes valeurs de τ_0 (seuil de 0,99)

Sur les 24 questions conservées :

Avec $\tau_0 = 0,10$, il ne reste aucun couple

Avec $\tau_0 = 0,20$, il ne reste aucun couple...

Avec $\tau_0 = 0,30$, il ne reste que trois couples qui feront l'objet d'une étude particulière.

Il faut attendre le niveau $\tau_0 = 0,40$ pour avoir un nombre de couples significatif, mais pratiquement aucune chaîne.

Ainsi, de la chaîne implicative envisagée (la seule apparue dans cette étude), il ne reste quasiment rien au niveau inclusif (voir graphe).

S'agissant d'une épreuve de concours et compte tenu des observations déjà faites sur les aspects dimensionnels ou plutôt non-dimensionnels des compétences mathématiques (Bodin, 1993), les observations qui précèdent doivent plutôt être portées au crédit de la qualité de l'épreuve. Tout se passe comme si l'épreuve permettait de contrôler 24 compétences quasi-indépendantes. Avec 24 questions il doit être difficile de faire mieux (mais pas impossible en réduisant encore le nombre de quasi-inclusions et, au delà, les tendances implicatives observées...).

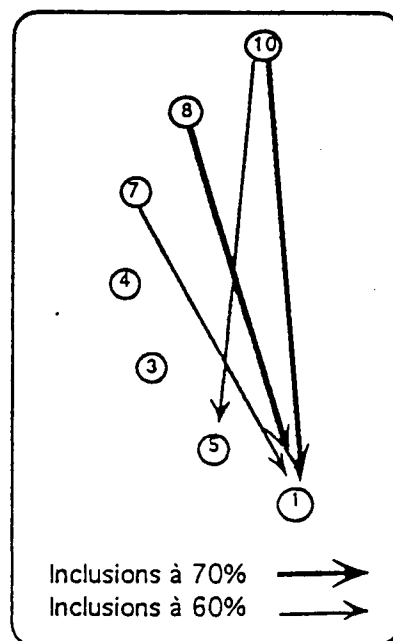
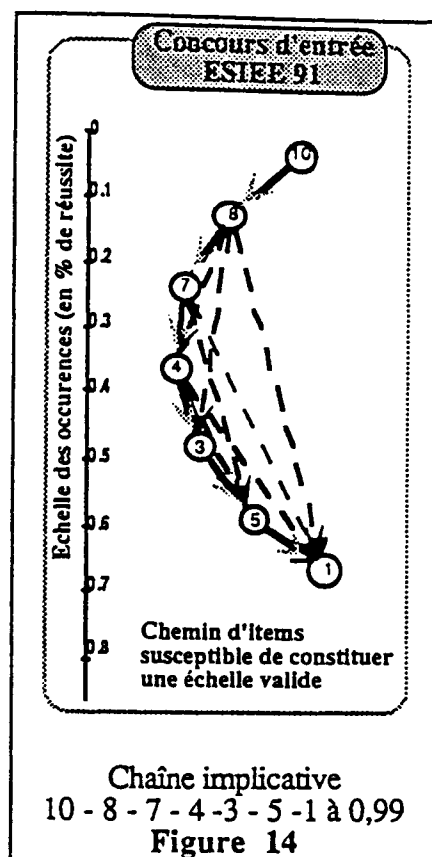


Figure 15

Exemple 2 : étude RADICAUX (EVAPM3/90)

Repris avec une épreuve davantage centrée sur un thème particulier (questions sur les radicaux d'EVAPM3/90), le nombre de liaisons "quasi-inclusives" est apparu comme beaucoup plus important.

Voici d'abord le graphe des implications statistiques présentant les fermetures transitives aux seuils 0,95 et 0,99.

Le test d'hypothèse inclusif étant fait pour $p = 0,10$, $p = 0,20$, $p = 0,3$, au seuil de confiance de 0,99, on observe que, par exemple :

La chaîne implicative 13, 10, 12, 11, peut être considérée comme constituant encore une chaîne inclusive à 80% (au seuil de confiance indiqué). De plus, cette chaîne est transitive.

Par contre elle ne constitue pas une chaîne inclusive à 90%.

Cependant, la chaîne 13, 10, 11 constitue une chaîne inclusive à 90%.

Si l'on considère le chemin implicatif

13, 08, 03, 01, 07

qui est transitif au seuil 0,95, les sous chemins inclusifs sont :

Inclusions à 80% : 13, 08, 03, 01, 07 et le chemin est transitif.

Inclusions à 90% : il reste encore les chaînes transitives 13, 08, 07 ; 13, 03, 01 et 13, 8, 1.

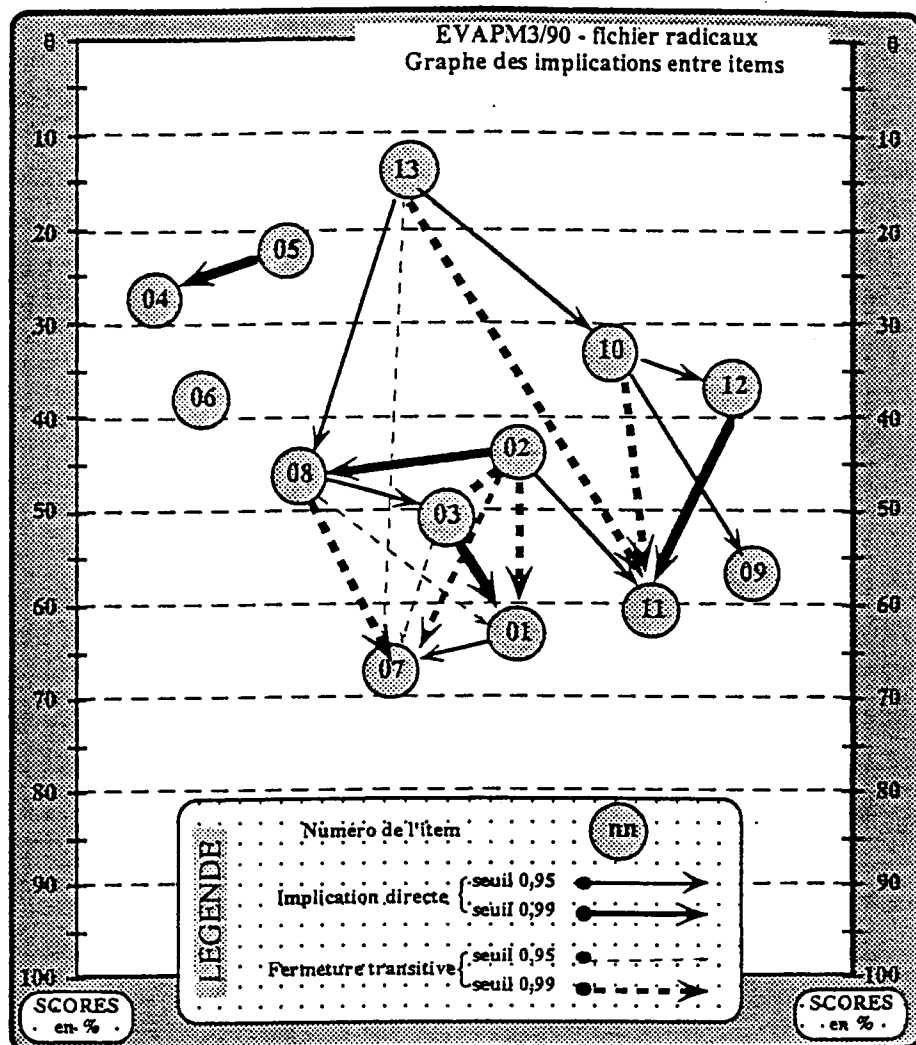


Figure 16

Lorsque l'on compare les résultats obtenus avec l'étude ESIEE et ceux obtenus ici, on est frappé de la différence, mais pas nécessairement étonnés. Avec l'étude *radicaux* on se trouve en présence de questions portant sur quelques compétences bien ciblées et en cours de développement... Il est alors souhaitables que ces questions présentent des gradations de difficulté et permettent de repérer le niveau de compétence dans un domaine bien défini.

En fait, l'analyse implicative complétée par notre test d'hypothèse et l'utilisation des modèles de réponse aux items (Bodin 96), se révèle être un bon outil pour l'analyse et la délimitation des champs conceptuels tout autant que pour l'étude de la validité des épreuves d'examens et de concours.

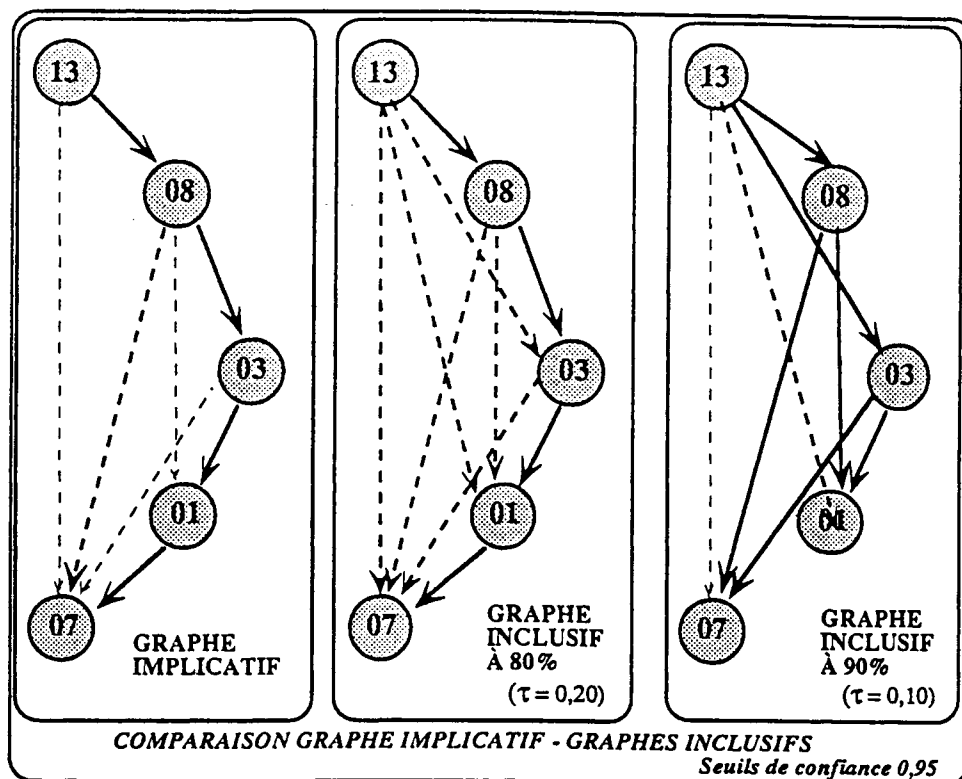


Figure 17

Relation entre quasi-inclusion on au seuil τ_0 et l'implication statistique au même seuil.

De par leurs construction il est clair que les énoncés $a \xrightarrow{\alpha} b$ et $\mathcal{A} \xrightarrow[1-\tau_0]{\alpha} \mathcal{B}$ ne sont pas synonymes.

Nous avons déjà vu que $a \xrightarrow{\alpha} b$ pouvait être vérifiée sans que $\mathcal{A} \xrightarrow[1-\tau_0]{\alpha} \mathcal{B}$ le soit, mais on

pourrait penser que $\mathcal{A} \xrightarrow[1-\tau_0]{\alpha} \mathcal{B}$ implique $a \xrightarrow{\alpha} b$ (implication logique).

L'exemple ci-contre montre qu'il n'est rien et que l'implication statistique et l'inclusion statistique doivent être considérées comme des notions distinctes.

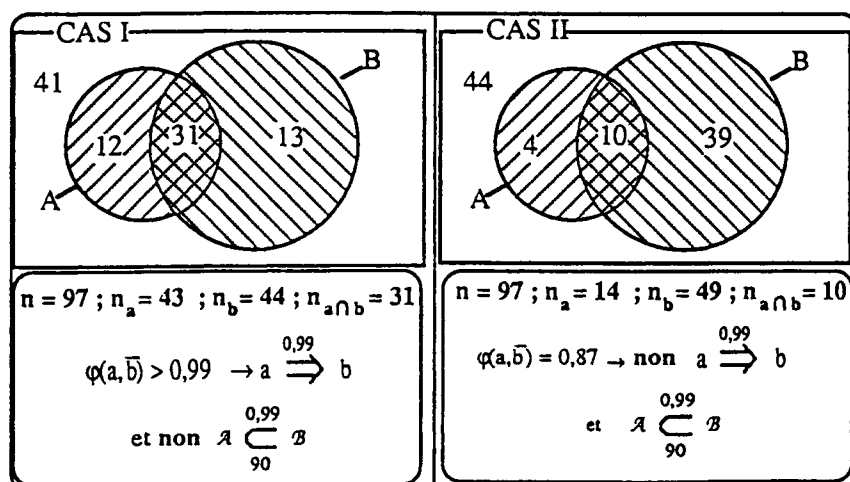


Figure 18

Références et bibliographie

- Ag Almouloud S. : 1992, *L'ordinateur: outil d'aide à l'apprentissage de la démonstration et de traitement d'analyse de données didactiques*, Thèse de l'Université de Rennes 1, Rennes.
- Bailleul M. : 1994, *Analyse statistique implicative : application à la modélisation de l'enseignement dans le système didactique*, Thèse, Université de Rennes 1, Rennes.
- Bodin A. & Couturier R. : 1996, *CHIC, Classification Hiérarchique et Cohésive, version sous windows, notice d'installation et d'utilisation*, ARDM/IRMAR Rennes
- Bodin A. : 1996, Improving the Diagnostic and Didactic Meaningfulness of Mathematics Assessment in France, *Annual Meeting of the American Educational Research Association AERA* - New-York
- Bodin A. : 1996, L'évaluation du savoir mathématique - Questions et méthodes (à paraître dans RDM).
- Bodin A.: 1993, What does to assess mean, *Investigations into Assessment in Mathematics Education, An ICMI Study* (ed Mogens NISS) - Kluwer Academic Publishers - Dordrecht
- Bodin A. & Couturier R. & Gras R. : 1996, *Analyse d'une épreuve de concours par la méthode implicative*. Communication aux journées de la société Française de Classification, Vannes
- Bodin A. : 1996, Mesures pour le système éducatif, *Actes des 7èmes Entretiens de la Villette*, 148-160, Centre National de Documentation pédagogique, Paris.
- Brousseau G.: 1993, *Stratégies de l'analyse statistique* . LADIST, Université de Bordeaux 1, Bordeaux.
- Gras R. : 1996, *L'implication statistique. Nouvelle méthode exploratoire de données*. La Pensée Sauvage. Grenoble
- Gras R. : 1979, *Contribution à l'étude expérimentale et à l'analyse de certaines acquisitions cognitives et de certains objectifs didactiques en mathématiques* - Thèse, Université de Rennes 1, Rennes.
- Gras R. : 1992, Data analysis : a method for the processing of didactic questions. In R. Douady & A. Mercier, (eds) *Research in Didactique of mathematics - selected papers* , La pensée Sauvage, Grenoble
- Gras R. : 1992, L'analyse des données: une méthodologie de traitement de questions de didactique, *Recherches en Didactique des Mathématiques, Vol. 12-1*, La pensée Sauvage, Grenoble.
- Gras R. : 1995, *Méthodes d'analyses statistiques multidimensionnelles en didactique des mathématiques*. Actes du colloque ARDM de CAEN (27 - 29 janvier 1995) - publié par l'ARDM, Orléans - Rennes
- Gras R. & PECAL M. : 1995, *L'évaluation en mathématiques : perspectives institutionnelles, pédagogiques et statistiques*. Actes de l'université d'été de l'APMEP - Sophia Antipolis 10-14 juillet 1995 - Brochure N° 102 de l'APMEP, Paris
- Larher A.: 1991, *Implication statistique et applications à l'analyse de démarches de preuve mathématique*, Thèse de l'Université de Rennes 1, Rennes.
- Larher A.: 1993, Analyses de similarités et d'implications entre procédures d'élèves dans de courtes démonstrations de géométrie, *Petit x*, n° 32, Grenoble.
- Polya G.: 1958, *Les mathématiques et le raisonnement plausible* , Gauthier Villars, Paris.
- Ratsimba-Rajohn H. : 1992, *Contribution à l'étude de la hiérarchie implicative. Application à l'analyse de la gestion didactique des phénomènes d'ostension et de contradiction* , Thèse de l'Université de Rennes 1, Rennes.
- Totohasina A.: 1992, *Méthode implicative en analyse de données et application à l'analyse de conceptions d'étudiants sur la notion de probabilité conditionnelle*, Thèse de l'Université de Rennes 1, Rennes.

Annexe 1

Eléments de tables pour la comparaison des indices d'implication
calculés suivant le modèle multinomial, binomial ou de Poisson -
calculs exacts ou approchés

CAS N = 50 ; B = 30 ; A = 20

card (A et non B)--	0	1	2	3	4	5	6	8	LEGENDE
Multinomial exact	1,00	1,00	1,00	1,00	0,98	0,93		0,38	
binomial exact	1,00	1,00	0,99	0,97	0,92	0,83		0,41	
binomial approché	1,00	1,00	0,99	0,97	0,94	0,88		0,50	
poisson exact	1,00	1,00	0,99	0,96	0,90	0,81		0,41	
Poisson approché	1,00	0,99	0,98	0,96	0,92	0,86		0,50	
				(**)	(***)		(*)		

(*) : Espérance mathématique sous
sous l'hypothèse d'indépendance
(**) : Valeurs critiques
au seuil de 0,99
(***) : Valeurs critiques
au seuil de 0,95

CAS N = 100 ; B = 60 ; A = 40

card (A et non B)--		3	5	7	8	9	10	11	12		16
Multinomial exact		1,00	1,00	1,00	1,00	1,00	0,99	0,97	0,93		0,42
binomial exact		1,00	1,00	0,99	0,99	0,97	0,94	0,89	0,83		0,43
binomial approché		1,00	1,00	0,99	0,99	0,97	0,95	0,91	0,86		0,50
poisson exact		1,00	1,00	0,99	0,98	0,96	0,92	0,87	0,81		0,43
poisson approché		1,00	1,00	0,99	0,98	0,96	0,93	0,89	0,84		0,50
					(**)	(**)	(**)	(*)			(*)

CAS N = 200 ; B = 120 ; A = 80

card (A et non B)--		18	19	20	21		22	23	24		32
Multinomial exact											
binomial exact											
binomial approché		1,00	0,99	0,99	0,98		0,97	0,96	0,94		0,50
poisson exact		0,99	0,99	0,98	0,97		0,96	0,94	0,91		0,45
poisson approché		0,99	0,99	0,98	0,97		0,96	0,94	0,92		0,50
				(**)			(**)				(*)

CAS N = 500 ; B = 300 ; A = 200

card (A et non B)--		59	60	61		65	66	67	...	80
Multinomial exact										
binomial exact										
binomial approché		0,99	0,99	0,99		0,97	0,96	0,94		0,50
poisson exact		0,99	0,99	0,98		0,95	0,94	0,92		0,47
poisson approché		0,99	0,99	0,98		0,95	0,94	0,93		0,50
				(**)		(***)				(*)

CAS N =1000 ; B = 600 ; A = 400

A et non B		130	131	132	133	134	135		139	140	141	142		160
Multinomial exact														
binomial exact														
binomial approché		1,00	0,99	0,99	0,99	0,99	0,98		0,96	0,96	0,95	0,94		0,50
poisson exact		0,99	0,99	0,99	0,98	0,98	0,98		0,95					
poisson approché		0,99	0,99	0,99	0,98	0,98	0,98		0,95	0,94	0,93	0,92		0,50
			(**)	(**)	(**)	(**)			(**)	(**)				(*)

A. BODIN - 6 janvier 96