



Calcul de l'aire sous le graphe de la fonction de répartition

Eléments théoriques et règles de calcul pratiques

Bernard Beauzamy

18/01/2020, rev. 28/07/2020, rev. 15/08/2026

Résumé opérationnel

Etant donnée une variable aléatoire réelle Y , définie au moyen d'une densité de probabilité, ou bien définie de manière empirique par une succession de valeurs enregistrées y_k , si M désigne la plus grande valeur que peut prendre Y , l'aire sous le graphe de Y est donné par la formule :

$$A = M - E(Y),$$

où $E(Y)$ désigne l'espérance de Y , si Y possède une densité, ou bien la moyenne arithmétique des valeurs prises, si Y est connue de manière empirique.

Etant donné un paramètre X , susceptible d'influer sur Y , nous montrons comment comparer les fonctions de répartition de Y , conditionnées par $X > \mu$ et $X < \mu$, μ étant la médiane de X . Vérifier la position respective de ces courbes est essentiel pour déterminer si X a une influence sur Y , car cette influence est déterminée par la différence des aires associées à chacune des deux fonction.

1. Notations

La fonction de répartition, notée $F(y)$, d'une variable aléatoire Y est définie par la formule :

$$F(y) = P(Y \leq y).$$

C'est une fonction croissante de y , à valeurs entre 0 et 1. L'aire sous le graphe joue un rôle essentiel dans notre méthode de "hiérarchisation de paramètres" : la variable Y est la variable d'intérêt ; étant donné un paramètre quelconque X , on conditionne Y par diverses contraintes de type $X < m$ ou $X > m$ et on compare les aires dans chaque cas. Selon l'importance de ces aires, X a plus ou moins d'influence sur Y . Il est donc important en pratique de pouvoir calculer cette aire rapidement et correctement.

2. Approche théorique

En théorie, le domaine de définition de F peut s'étendre de $-\infty$ à $+\infty$, mais en pratique il s'agit toujours d'un intervalle borné $[m, M]$. On a $F(m) = 0$, $F(M) = 1$.

La fonction F est la primitive de la densité f , nulle en m . L'aire sous le graphe de F vaut :

$$A = \int_m^M F(y) dy.$$

En intégrant par parties, et en notant $E(Y)$ l'espérance de Y :

$$A = \int_m^M F(y) dy = [yF(y)]_m^M - \int_m^M yF'(y) dy = M - \int_m^M yf(y) dy = M - E(Y).$$

Cette formule très simple demeure correcte lorsque la loi de Y n'est pas connue : on dispose seulement de relevés expérimentaux.

3. Cas discret : relevés expérimentaux

On a fait N expériences et on a observé les valeurs y_1 avec multiplicité $n_1, \dots, y_K$ avec multiplicité n_K , avec bien sûr $n_1 + \dots + n_K = N$; avec cette notation, K est le nombre de valeurs distinctes. On peut bien sûr supposer que les y_k sont rangés en ordre croissant :

$$m = y_1 < \dots < y_K = M.$$

Pour chaque valeur de y , la probabilité empirique $P(Y \leq y)$ est définie comme étant le nombre de fois où, dans la liste, une valeur $\leq y$ a été rencontrée, divisé par le nombre total d'expériences, noté ici N .

La plus petite valeur rencontrée dans la liste est $y_1 = m$; si $y < y_1$, on ne rencontre jamais de valeur inférieure à y , donc $F(y) = 0$ pour tout $y < y_1$.

En $y = y_1$, la fonction F prend la valeur $\frac{n_1}{N}$ et de même pour tout y , $y_1 \leq y < y_2$.

De même, $F(y) = \frac{n_1 + n_2}{N}$ si $y_2 \leq y < y_3$ et, plus généralement, pour tout $k = 1, \dots, K-1$:

$$F(y) = \frac{n_1 + \dots + n_k}{N} \text{ si } y_k \leq y < y_{k+1}, \quad F(y) = \frac{n_1 + \dots + n_{K-1}}{N} \text{ si } y_{K-1} \leq y < y_K = M,$$

et finalement :

$$F(y) = 1 \text{ si } y \geq y_K.$$

La fonction F est donc constante sur une succession d'intervalles, fermés à gauche, ouverts à droite. Les valeurs prises vont en croissant. Le fait que les intervalles de définition soient fermés ou ouverts à chaque extrémité est sans importance du point de vue de l'aire.

Le $k^{\text{ème}}$ intervalle, noté I_k , va de y_k à y_{k+1} pour $k = 1, \dots, K-1$; le dernier intervalle, I_{K-1} , va de y_{K-1} à M .

La taille du $k^{\text{ème}}$ intervalle est $y_{k+1} - y_k$ pour $k = 1, \dots, K-1$.

L'aire sous le graphe de F est donc :

$$A = \sum_{k=1}^{K-1} \frac{n_1 + \dots + n_k}{N} (y_{k+1} - y_k) + M - y_K,$$

Ce qu'on écrit :

$$\begin{aligned} A &= \frac{n_1}{N} (y_2 - y_1) + \frac{n_1 + n_2}{N} (y_3 - y_2) + \dots + \frac{n_1 + \dots + n_k}{N} (y_{k+1} - y_k) + \dots + \frac{n_1 + \dots + n_{K-1}}{N} (y_K - y_{K-1}) + M - y_K \\ &= -\frac{n_1}{N} y_1 - \frac{n_2}{N} y_2 - \dots - \frac{n_k}{N} y_k - \frac{n_{K-1}}{N} y_{K-1} - \frac{n_K}{N} y_K + M \end{aligned}$$

et donc :

$$A = M - \frac{1}{N} \sum_{k=1}^K n_k y_k$$

La quantité $\frac{1}{N} \sum_{k=1}^K n_k y_k$ peut être considérée comme l'espérance expérimentale de Y : c'est celle qui résulte de l'échantillon de mesure. La formule $A = M - E(Y)$ demeure donc valable.

La valeur $E(Y)$ est simplement la moyenne de tous les résultats expérimentaux de Y ; il n'est pas nécessaire de les trier par ordre croissant pour la calculer.

4. Conditionnement de Y

Pour l'application de la méthode de hiérarchisation de paramètres, on est amené à conditionner Y par des situations du type $X < \mu$ ou $X > \mu$, où X est un paramètre quelconque. Les domaines de variation $[m, M]$ pour Y sont différents dans les deux cas ; notons-les respectivement $[m_1, M_1], [m_2, M_2]$.

L'application des règles précédentes est alors très simple. Par exemple, pour la situation $X < \mu$, on va extraire du tableau initial toutes les lignes pour lesquelles $X < \mu$; soit T_1 le tableau de données ainsi réduit. On va calculer la moyenne des valeurs de Y apparaissant dans ce tableau réduit ; notons-la $E_1 = E(Y | X < \mu)$; l'aire au-dessous du graphe de la fonction de répartition dans la situation $X < \mu$ sera :

$$A_1 = M_1 - E_1.$$

On fait de même avec la situation $X > \mu$; on extrait le tableau T_2 et on calcule la moyenne $E_2 = E(Y | X > \mu)$ des valeurs de Y dans ce tableau. L'aire au-dessous du graphe de la fonction de répartition dans la situation $X > \mu$ sera :

$$A_2 = M_2 - E_2.$$

Remarque. – Pour l'utilisation de la méthode de hiérarchisation, il est recommandé de tracer les deux courbes $F_1(y) = P(Y \leq y | X < \mu)$ et $F_2(y) = P(Y \leq y | X > \mu)$ de manière à vérifier que l'une est constamment au-dessous de l'autre. Ceci peut être fait simplement, comme nous allons maintenant le voir.

5. Méthode pour déterminer simplement le positionnement relatif des courbes

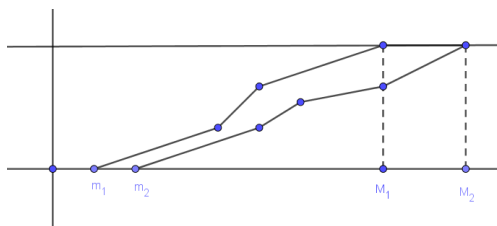
On reprend le tableau de tous les résultats de Y . On le range par ordre de y croissants. On a deux sous-tableaux, l'un pour $X < \mu$ (lignes rouges), l'autre pour $X > \mu$ (lignes bleues) ; évidemment, une ligne ne peut être à la fois rouge et bleue. On élimine celles qui ne sont ni rouges, ni bleues. Soit N_1 le nombre de lignes rouges et N_2 le nombre de lignes bleues.

On dispose alors d'un tableau à deux colonnes : en première colonne, la valeur de Y (croissante au sens large) ; en deuxième colonne l'indication "rouge" ou "bleu".

On ajoute trois colonnes :

- Colonne 3 : on parcourt le tableau en commençant par la première ligne. A chaque fois que l'on rencontre l'indication "rouge" sur cette ligne, on incrémente de $\frac{1}{N_1}$: on met $1/N_1$ pour la première fois que l'on rencontre "rouge", $2/N_1$ la seconde, etc.
- Colonne 4 : la même chose, pour l'indication "bleu", et en mettant $1/N_2$ au lieu de $1/N_1$.
- Colonne 5 : la différence entre colonne 3 et colonne 4. Cette différence doit garder un signe constant pour que la méthode de hiérarchisation puisse être appliquée.

Cette vérification peut être automatisée.



Remarque. - Si l'on veut que $F_1 \geq F_2$, il faut évidemment que $m_1 \leq m_2$ et que $M_1 \leq M_2$; voir figure.