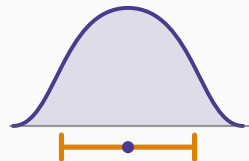


Statistique inférentielle

Chapitre 17 : Comparer deux groupes



Jaouad Madkour

Faculté d'Économie et de Gestion de Tanger · 27 août 2026

Comparer deux moyennes

Comparer deux proportions

Les échantillons appariés

Ce qu'il faut retenir

Comparer deux moyennes

On teste l'écart, pas les moyennes

$$H_0 : m_1 = m_2 \quad \text{c'est-à-dire} \quad H_0 : m_1 - m_2 = 0$$

La statistique porte sur la **différence**, qui est elle-même une variable aléatoire.

L'écart-type de l'écart

Les deux groupes étant indépendants, les variances s'additionnent — c'est le chapitre 12 des probabilités :

$$\sigma(\bar{X}_1 - \bar{X}_2) = \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}$$

Le test de Welch

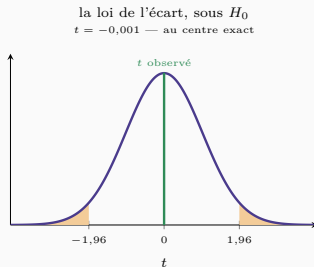
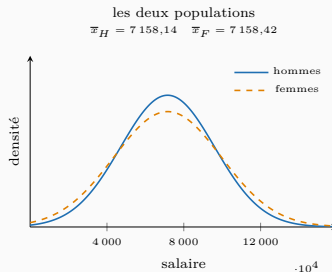
$$t = \frac{\bar{x}_1 - \bar{x}_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}}$$

C'est la version qui **ne suppose pas** l'égalité des variances. C'est celle qu'appliquent R et Python par défaut, et il n'y a aucune raison d'en préférer une autre.

Une simplification qui n'en est pas une

Les manuels présentent souvent une version « à variances égales », dont le degré de liberté est simplement $n_1 + n_2 - 2$.

Elle exige une hypothèse supplémentaire, qu'il faudrait tester au préalable — et elle ne gagne presque rien. **Welch d'abord, toujours.**



Un écart de 0,28 dirham sur des salaires de sept mille : le test tombe au centre exact.

$$t = \frac{7\,158,14 - 7\,158,42}{239,55} = -0,001, \text{ soit } p \approx 1. \text{ On ne rejette pas } H_0.$$

Et pourtant on n'a rien prouvé. Ne pas rejeter, c'est constater l'absence de preuve d'un écart — non prouver l'égalité. Avec ces effectifs, un écart réel de 300 dirhams aurait été détecté ; un écart de 50 serait passé inaperçu.

La question utile n'est donc pas « est-ce significatif ? » mais « quel écart mon échantillon était-il capable de détecter ? » — c'est la puissance, chapitre 14.

Hommes contre femmes, sur la population entière

Groupe	n	Moyenne	Écart-type
Hommes	285	7 158,14	2 448,48
Femmes	215	7 158,42	2 795,58

$$\sqrt{\frac{2\,448,48^2}{285} + \frac{2\,795,58^2}{215}} = 239,55$$

$$t = \frac{-0,28}{239,55} = -0,001 \quad p \approx 1$$

Le test tombe au centre exact

Un écart de vingt-huit centimes sur des salaires de sept mille dirhams. Il est difficile d'imaginer un résultat plus clairement non significatif.

Et pourtant **on n'a rien prouvé**. On a constaté qu'aucun écart n'est détectable — ce qui est une information, mais pas une démonstration d'égalité.

Ce qu'il faut ajouter pour que la conclusion serve

Avec ces effectifs, l'erreur-type de l'écart vaut 240 dirhams. Le test détecte donc facilement un écart de 500, difficilement un écart de 200, jamais un écart de 50.

C'est cela qu'il faut écrire dans le rapport, et non simplement « pas de différence significative ».

Comparer deux proportions

Le test de deux proportions

$$z = \frac{f_1 - f_2}{\sqrt{\hat{p}(1 - \hat{p}) \left(\frac{1}{n_1} + \frac{1}{n_2} \right)}} \quad \text{avec} \quad \hat{p} = \frac{k_1 + k_2}{n_1 + n_2}$$

Pourquoi une proportion commune

Sous H_0 , les deux groupes ont **la même** proportion. La meilleure estimation de cette valeur unique utilise donc les deux échantillons réunis.

C'est la même logique qu'au chapitre 16, où l'on employait p_0 et non f : **sous H_0 , on calcule comme si H_0 était vraie.**

Deux agences, deux taux de conversion

Agence A : 48 ventes sur 300 visites, soit 16,0 %. Agence B : 72 ventes sur 360 visites, soit 20,0 %.

$$\hat{p} = \frac{120}{660} = 0,1818 \quad z = \frac{-0,04}{\sqrt{0,1818 \times 0,8182 \times \left(\frac{1}{300} + \frac{1}{360}\right)}} = -1,32$$

$|-1,32| < 1,96$: on ne rejette pas, $p = 0,186$.

La décision, malgré l'absence de significativité

Quatre points d'écart de conversion représentent des sommes considérables sur l'année.

Le test dit seulement qu'avec six cent soixante visites, on ne peut pas exclure le hasard. **La bonne décision n'est pas d'ignorer l'écart : c'est de prolonger l'observation jusqu'à pouvoir trancher.**

Les échantillons appariés

Un cas particulier fréquent

Quand les deux mesures portent sur les mêmes individus

Avant et après une formation, avant et après un changement de tarif, deux méthodes sur les mêmes clients.

On ne compare **pas** deux groupes : on calcule les n **différences individuelles**, et l'on teste si leur moyenne est nulle.

C'est le test du chapitre 16, appliqué aux écarts

$$t = \frac{\bar{d}}{s_d/\sqrt{n}} \sim \text{Student à } n - 1 \text{ ddl}$$

Aucune formule nouvelle : on a transformé un problème à deux échantillons en un problème à un seul.

L'appariement élimine la variabilité individuelle

Traiter deux mesures sur les mêmes personnes comme deux groupes indépendants gaspille l'information la plus précieuse : le fait que ce soient **les mêmes personnes**.

Les différences individuelles sont beaucoup moins dispersées que les niveaux. À écart réel identique, un test apparié détecte souvent avec dix fois moins d'observations.

La règle de reconnaissance

Si l'on peut **mettre les observations deux par deux**, il faut appairer.

C'est la faute la plus coûteuse de ce chapitre : conduire une étude avant-après, puis l'analyser comme deux groupes indépendants, et conclure à tort qu'il n'y a pas d'effet.

Ce qu'il faut retenir

Six points

1. On teste la **différence** : $H_0 : m_1 - m_2 = 0$.
2. Les variances s'additionnent : $\sigma(\bar{X}_1 - \bar{X}_2) = \sqrt{s_1^2/n_1 + s_2^2/n_2}$.
3. Utiliser **Welch** par défaut : il ne suppose pas l'égalité des variances.
4. Sur nos salaires, $t = -0,001$: aucun écart détectable — ce qui ne prouve pas l'égalité.
5. Deux proportions : on estime une proportion **commune** sous H_0 .
6. Si les observations vont **deux par deux**, il faut appairer : le gain de puissance est considérable.

La suite

Trois groupes ou davantage, et la méthode ne tient plus : comparer les diplômes deux à deux ferait trois tests, donc un risque réel bien supérieur à 5 %.

Le chapitre 18 introduit l'outil conçu pour cela — et l'on y retrouvera la décomposition de la variance du cours de statistique descriptive.