

Analyse nutritionnelle : PCIMA et SMART

Anthropométrie, indicateurs d'admission et de sortie, et une enquête SMART analysée de bout en bout au regard des normes de croissance de l'OMS et des seuils de performance Sphere.

Formateur	Pierre Robentz Cassion
Niveau	Intermédiaire
Heures apprenant	24 h
Enseignée en	Python · R
Mise à jour	28 juillet 2026
Secteurs	Nutrition
Normes et méthodologies	Normes OMS de croissance de l'enfant · Enquête SMART · Standards Sphère · Cadre intégré de classification de la sécurité alimentaire (IPC) · Définitions d'indicateurs de l'UNICEF

Consulter et exécuter la formation en ligne sur data-analysis.cassion.dev/fr/cours/nutrition-analysis-cmam-smart

Le contenu pédagogique est diffusé sous licence CC BY 4.0. Vous pouvez l'adapter et le rediffuser, y compris à des fins commerciales, en créditant Cassion.

Ce que vous saurez faire

- Calculer les scores z poids-taille, taille-âge et poids-âge selon les normes OMS 2006, à partir de mesures brutes
- Appliquer correctement les définitions de cas de la MAS et de la MAM, y compris la règle des œdèmes qui prime sur l'anthropométrie
- Dire où le PB et le poids-pour-taille divergent, de combien, et quels enfants chaque critère trouve
- Produire un rapport de plausibilité SMART et décider si une enquête doit être acceptée
- Estimer les prévalences de la MAG et de la MAS avec un intervalle ajusté du plan et les lire au regard des phases IPC de malnutrition aiguë
- Calculer la performance PCIMA sur le dénominateur que Sphere spécifie réellement, et trouver le site que le chiffre du district masque
- Expliquer pourquoi le rapport admissions sur charge attendue n'est pas la couverture, et ce qu'exige son estimation

Prérequis

Aucun. Cette formation ne suppose aucune expérience préalable de la programmation.

Sommaire

I	Mesurer un enfant	1
1	Les quatre mesures, et à quoi sert chacune	2
1.1	Quatre choses, mesurées sur un enfant	2
1.2	Longueur ou taille, et les 0,7 cm qui ne sont pas un arrondi	2
1.3	Le PB, et pourquoi on l'emploie	3
1.4	La précision décide de ce qu'une mesure peut soutenir	3
1.5	Lisez la distribution avant tout calcul	4
1.6	Ce qu'un analyste nutrition vérifie d'abord	4
1.7	La suite	5
2	Du poids et de la taille au score z	6
2.1	Ce que dit un score z	6
2.2	Trois scores z, trois questions	6
2.3	La méthode LMS	6
2.4	La clé de recherche a trois parties	7
2.5	Le calcul	7
2.6	Le signalement — deux règles, deux réponses	8
2.7	L'écart-type est déjà un constat	8
2.8	Enregistrez les scores z avec leurs entrées	9
2.9	La suite	9
II	Qui est un cas	10
3	Les définitions de cas, et la clause que tout le monde oublie	11
3.1	Les définitions	11
3.2	L'implémentation fausse	11
3.3	L'implémentation juste	12
3.4	Vérifiez que les parties totalisent le tout	12
3.5	Les définitions par le PB, sur le registre	13
3.6	Ce qu'une définition ne dit pas	13
3.7	La suite	14
4	Là où les deux critères divergent	15
4.1	La question à laquelle le registre peut répondre	15
4.2	Calculez les deux sur les mêmes enfants	15
4.3	Le tableau croisé	15
4.4	Le désaccord, c'est l'âge	16
4.5	Lequel a raison ?	17
4.6	Ce que cela implique pour vos chiffres	17
4.7	Rapportez toujours la paire	18
4.8	La suite	18

III	Juger une enquête	19
5	Cette enquête est-elle croyable ?	20
5.1	Une enquête est acceptée ou ne l'est pas	20
5.2	Les contrôles	20
5.3	Les signalements	20
5.4	L'écart-type	21
5.5	La préférence de chiffres	21
5.6	L'attraction des âges	22
5.7	Sex-ratio et biais d'équipe	22
5.8	Le verdict	23
5.9	Rédigez sans accuser personne	24
5.10	La suite	24
6	L'intervalle qui chevauche deux phases IPC	25
6.1	Le nombre pour lequel l'enquête existe	25
6.2	L'estimation, avec le plan	25
6.3	Les phases IPC	26
6.4	Ce que cela veut dire, dit franchement	26
6.5	Ce qui se serait passé sans l'effet de plan	27
6.6	Lisez-le avec le rapport de plausibilité	27
6.7	Le bloc de rapportage complet	28
6.8	La suite	28
IV	Juger un programme	29
7	Le dénominateur qui déplace le taux de guérison de neuf points	30
7.1	Quatre issues et un standard	30
7.2	Les trois dénominateurs candidats	30
7.3	Celui que Sphere désigne	31
7.4	Décomposez par site	31
7.5	La cause est en général la distance, non l'observance	32
7.6	Le gain de poids, et les saisies qui ne peuvent pas être justes	33
7.7	Le tableau de performance à publier	33
7.8	La suite	34
8	Pourquoi admissions sur charge attendue n'est pas la couverture	35
8.1	Le chiffre que tout le monde rapporte	35
8.2	Ce qu'il mesure réellement	35
8.3	Ce qu'exige la couverture	36
8.4	La question des barrières est la moitié utile	37
8.5	Quoi rapporter sans enquête de couverture	37
8.6	Ce que ce cours vous laisse	38

À propos de cette formation

C'est le premier cours du programme où le contenu sectoriel est la partie difficile. Tout ce qui relevait des modules 2 et 3 était de la méthode transférable d'un secteur à l'autre. Ce n'est pas le cas ici — les normes de croissance de l'OMS, les critères de plausibilité SMART, les seuils de performance Sphere et les phases IPC sont précis, publiés, lourds de conséquences, et ce public est tenu aux quatre.

Le cours repose sur trois registres qui couvrent ensemble tout le cycle de programme. Une **enquête SMART** de 930 enfants, livrée en poids et taille bruts pour que les scores z soient calculés. Un **registre de dépistage** de 4 218 enfants, là où les cas sont trouvés. Et un **registre d'admissions PCIMA** de 1 100 épisodes, là où ils sont traités et où la performance se juge.

Trois constats tirés de ces données donnent sa forme au cours, et chacun est un nombre réel calculé sur les fichiers livrés, non une affirmation.

Les deux critères d'admission trouvent des enfants différents. Parmi les enfants sévères selon l'une ou l'autre mesure, 468 le sont selon le seul poids-pour-taille, 24 selon le seul PB et 142 selon les deux — une concordance de 22,4 %. Le désaccord, c'est l'âge : les enfants trouvés par le seul PB ont 14,3 mois en moyenne, ceux trouvés par le seul poids-pour-taille en ont 31,4.

Le dénominateur décide du taux de guérison. Calculé comme Sphere le spécifie — sur les enfants ayant atteint une issue, transferts exclus — le programme guérit 82,3 %. Divisez le même numérateur par les 1 100 admissions et il affiche 73,6 %.

Le chiffre du district masque le site. L'abandon est de 13,7 % au total, sous le maximum Sphere de 15 %, et de 20,5 % sur un site, au-dessus.

Python et R tout du long. Vous devriez avoir suivi le module 3, ou savoir au moins défendre un dénominateur et mettre un intervalle sur une proportion — ce cours ajoute les définitions cliniques aux deux.

Première partie

Mesurer un enfant

CHAPITRE 1

Les quatre mesures, et à quoi sert chacune

Leçon 1 sur 8 · 120 min

Poids, longueur ou taille, PB et œdèmes. Chacune répond à une question différente, chacune a une précision qui décide de ce qu'elle peut soutenir, et celle qui n'est pas du tout une mesure est celle qui prime sur les autres.

1.1 Quatre choses, mesurées sur un enfant

Tout ce cours repose sur quatre observations, et un analyste qui ignore comment chacune est prise les lira toutes de travers.

Mesure	Unité	Précision	Ce à quoi elle est sensible
Poids	kg	0,1 kg	Déficit aigu — il chute vite et se récupère vite
Longueur / taille	cm	0,1 cm	Déficit chronique — il s'accumule et ne se récupère pas
PB	mm	1 mm	Déficit aigu, et le risque de mortalité plus directement que le
Œdèmes	présents / absents	—	Pas une mesure ; un signe clinique, et il prime

Deux de ces lignes méritent qu'on s'y arrête, car elles expliquent l'essentiel de ce qui suit.

Le poids bouge, la taille non. Un enfant qui a eu faim deux mois est plus léger qu'un enfant de même taille qui n'a pas eu faim, et il reprendra ce poids en quelques semaines s'il est nourri. Un enfant qui a eu faim deux ans est *plus petit*, et il ne rattrapera pas cette taille. Ce seul fait explique pourquoi le secteur distingue l'émaciation du retard de croissance et pourquoi les deux exigent des indicateurs et des programmes différents.

Les œdèmes ne sont pas sur une échelle. C'est un signe clinique — œdèmes bilatéraux prenant le godet, vérifiés en pressant les deux pieds trois secondes — et un enfant qui en présente est en malnutrition aiguë sévère quels que soient son poids et son périmètre brachial. Toute définition de cas de l'unité suivante porte une clause d'œdème, et toute analyse qui calcule la MAS comme « le décompte sous le seuil de score z » est fautive d'autant d'enfants œdémateux que le registre en contient.

1.2 Longueur ou taille, et les 0,7 cm qui ne sont pas un arrondi

Les enfants de moins de 24 mois sont mesurés **couchés** — c'est la longueur. Ceux de 24 mois et plus sont mesurés **debout** — c'est la taille. Le même enfant mesure environ 0,7 cm de plus couché que debout, car la colonne se décomprime.

Les normes de l'OMS sont publiées pour les deux, et la règle est la suivante.

- Sous 24 mois, employez la norme de **longueur**, et si l'enfant a été mesuré debout, **ajoutez 0,7 cm** avant la recherche.
- À 24 mois et plus, employez la norme de **taille**, et si l'enfant a été mesuré couché, **retranchez 0,7 cm**.

```
import pandas as pd
```

```
smart = pd.read_csv("smart-nutrition-survey-2024.v1.csv")
print(smart["measured_lying"].value_counts())
```

```
library(dplyr)
smart |> count(measured_lying)
```

752 mesurés couchés, 178 debout. **Le registre enregistre la position, ce qui rend l'ajustement possible** — et un registre qui ne l'enregistre pas est un registre dont les scores z portent une erreur non quantifiable pour chaque enfant mesuré dans la position non standard.

Inversez le signe et vous biaisez la moitié la plus jeune de l'enquête. Le laboratoire du cours sur les jointures vous l'a déjà fait faire ; voici pourquoi cela compte cliniquement.

1.3 Le PB, et pourquoi on l'emploie

Le PB est un ruban autour du bras, lu au millimètre, et il paraît fruste à côté d'un score z calculé à partir de deux instruments étalonnés. On l'emploie pour trois raisons qui décident ensemble de la plupart des programmes communautaires.

- **Il exige une mesure, pas deux.** Un agent de santé communautaire muni d'un ruban peut dépister un village ; un dépistage par poids-pour-taille exige une balance, une toise et une table de référence.
- **Il prédit la mortalité au moins aussi bien.** À âge donné, un bras petit prédit fortement le décès, et dans plusieurs études mieux que le poids-pour-taille.
- **Il n'exige aucun âge.** Ce qui compte énormément dans les populations sans actes de naissance, où l'âge dont un score z a besoin est lui-même une estimation.

Le coût fait l'objet de la leçon 4 — le PB croît avec l'âge indépendamment de l'état nutritionnel, si bien qu'un seuil fixe trouve des enfants différents selon l'âge.

1.4 La précision décide de ce qu'une mesure peut soutenir

```
muac = pd.read_csv("muac-screening-artibonite-2024.v1.csv")
values = muac.loc[muac["muac_mm"].notna(), "muac_mm"]

near_cutoff = values.between(123, 127).sum()
print(f"{near_cutoff} of {len(values)} children within 2 mm of the 125 mm cut-off")
```

```
muac |>
  filter(!is.na(muac_mm)) |>
  summarise(near = sum(between(muac_mm, 123, 127)), n = n())
```

Le PB se lit au millimètre et deux agents formés obtiennent couramment des mesures répétées différant de 2 à 3 mm. Tout enfant situé à environ 3 mm d'un seuil est donc un enfant dont la classification dépend de qui tenait le ruban.

Cela ne rend pas le seuil faux — un seuil doit bien être quelque part — mais cela a deux conséquences qu'un analyste doit porter.

- **Ne rapportez jamais un taux fondé sur le PB à deux décimales.** La mesure ne le soutient pas.

- **Attendez-vous à un amoncellement au seuil dans les données de programme.** Un agent qui lit 114 et sait que 115 est le seuil d'admission pose un jugement clinique, il ne falsifie pas un registre, et la distribution le montrera.

Le même argument vaut pour le poids au 0,1 kg et la taille au 0,1 cm, et la leçon 5 en fait un contrôle formel.

1.5 Lisez la distribution avant tout calcul

```
cmam = pd.read_csv("cmam-admissions-2024.v1.csv")

print(cmam[["muac_admission_mm", "weight_admission_kg", "height_cm"]].describe().round(1))
print(f"height missing: {cmam['height_cm'].isna().sum()} of {len(cmam)}")

cmam <- readr::read_csv("cmam-admissions-2024.v1.csv")

summary(select(cmam, muac_admission_mm, weight_admission_kg, height_cm))
sum(is.na(cmam$height_cm))
```

Soixante-douze des 1 100 admissions n'ont pas de taille. Ce n'est pas une gêne de qualité à nettoyer — c'est la raison pour laquelle le poids-pour-taille est incalculable pour ces enfants alors que le PB l'est, et donc la raison pour laquelle toute comparaison des deux critères a deux dénominateurs différents. La leçon 4 est bâtie exactement là-dessus.

1.6 Ce qu'un analyste nutrition vérifie d'abord

Cinq contrôles, dans l'ordre, avant tout indicateur.

- **La position est-elle enregistrée ?** Sans `measured_lying`, l'ajustement de 0,7 cm est une supposition.
- **L'âge est-il présent, et comment a-t-il été obtenu ?** Un score `z` exige l'âge ; une distribution d'âges attirée par les valeurs rondes dit qu'il a été estimé, ce que la leçon 5 mesure.
- **Les œdèmes sont-ils un champ distinct ?** S'ils sont fondus dans un indicateur général de « complications », le décompte de la MAS ne peut pas être calculé correctement.
- **Les unités.** PB en mm ou en cm, poids en kg, taille en cm. Le cours de nettoyage a trouvé sept saisies en centimètres dans le registre de dépistage ; elles sont toujours là.
- **Les plages plausibles.** PB 80-220 mm, poids 2-30 kg, taille 45-125 cm pour 6-59 mois. Au-delà, c'est une erreur de saisie, non un résultat.

```
implausible = cmam[
  ~cmam["muac_admission_mm"].between(80, 220)
  | ~cmam["weight_admission_kg"].between(2, 30)
  | ~cmam["height_cm"].between(45, 125).fillna(True)
]
print(f"{len(implausible)} admissions outside plausible ranges")

cmam |>
  filter(!between(muac_admission_mm, 80, 220) |
```

```
!between(weight_admission_kg, 2, 30) |  
(!is.na(height_cm) & !between(height_cm, 45, 125))) |>  
nrow()
```

L'anthropométrie est la seule partie des données de ce secteur où l'erreur de mesure est bien caractérisée et publiée. Servez-vous-en : elle dit quelle précision votre indicateur peut porter, et c'est la raison pour laquelle un écart de deux points entre deux enquêtes n'est en général rien.

1.7 La suite

Vous avez quatre mesures et savez à quoi sert chacune. La leçon suivante en transforme deux en la quantité sur laquelle repose toute prévalence de ce cours — un score z selon les normes OMS 2006, calculé à partir de la table de référence plutôt que lu dans une colonne de confiance.

CHAPITRE 2

Du poids et de la taille au score z

Leçon 2 sur 8 · 150 min

La méthode LMS, la recherche dans la table de référence, les trois scores z et ce que chacun mesure. Calculé sur 930 enfants, il donne une moyenne de -0,67 et un écart-type de 1,22 — et ce second nombre est déjà un constat.

2.1 Ce que dit un score z

Un score z répond à une seule question — **à quelle distance cet enfant est-il de la médiane d'une population de référence en bonne santé, en écarts-types ?**

Un score z poids-pour-taille de -2 signifie que cet enfant pèse ce que pèse un enfant situé à deux écarts-types sous la médiane, pour sa taille. Ce n'est un pourcentage de rien, et ce n'est pas comparable à un percentile sans conversion.

La référence est constituée par les **normes OMS 2006 de croissance de l'enfant**, bâties sur une étude multi-pays d'enfants élevés dans des conditions ne limitant pas la croissance — allaités, foyers non-fumeurs, soins adéquats. Cette construction est délibérée : les normes décrivent comment les enfants *devraient* grandir, non comment les enfants d'un lieu donné grandissent.

2.2 Trois scores z, trois questions

Score	Compare	Déficit détecté	Récupérable ?
Poids-pour-taille (WHZ)	Le poids à la taille	Aigu — émaciation	Oui, en semaines
Taille-pour-âge (HAZ)	La taille à l'âge	Chronique — retard de croissance	Non
Poids-pour-âge (WAZ)	Le poids à l'âge	Les deux à la fois — insuffisance pondérale	En partie

Le poids-pour-âge est celui dont il faut se méfier. Il mélange les deux déficits, si bien qu'un WAZ bas ne dit pas si l'enfant est émacié, en retard de croissance ou les deux, et la réponse programmatique diffère. Il survit en suivi de la croissance parce qu'il n'exige pas de toise, et il ne doit pas servir à classer la malnutrition aiguë.

Ce cours calcule le poids-pour-taille, car c'est sur lui que reposent les définitions de cas et les seuils IPC.

2.3 La méthode LMS

Les normes sont publiées sous forme de trois paramètres, par sexe, par mesure et par tranche de 0,1 cm de longueur ou de taille.

- **L** — la puissance de Box-Cox qui normalise une distribution asymétrique
- **M** — la médiane
- **S** — le coefficient de variation

$$z = ((\text{weight} / M)^L - 1) / (L * S)$$

```
import pandas as pd

reference = pd.read_csv("who-2006-weight-for-lenhei.csv")
print(reference.head(3))
print(f"{len(reference)} rows: sex x standard x lenhei in 0.1 cm steps")

reference <- readr::read_csv("who-2006-weight-for-lenhei.csv")
nrow(reference)
```

2 404 lignes. **Lisez la table; ne réimplémentez pas l'interpolation d'après un manuel.** R dispose du paquet officiel `anthro`; Python n'a pas d'équivalent maintenu qui s'installe proprement, et c'est exactement pourquoi cette table est livrée ici.

2.4 La clé de recherche a trois parties

C'est la partie qui dérape, et le laboratoire du cours sur les jointures était bâti dessus.

```
def lookup_key(row):
    lying = row["measured_lying"] == "true"
    standard = "L" if row["age_months"] < 24 else "H"

    lenhei = row["height_cm"]
    if standard == "L" and not lying:
        lenhei += 0.7
    elif standard == "H" and lying:
        lenhei -= 0.7

    return row["sex"], standard, round(lenhei, 1)

lookup_key <- function(sex, age, height, lying) {
  standard <- if (age < 24) "L" else "H"
  lenhei <- height +
    if (standard == "L" && !lying) 0.7 else if (standard == "H" && lying) -0.7 else 0
  list(sex = sex, lorh = standard, lenhei = round(lenhei, 1))
}
```

Trois éléments de cette fonction sont des décisions, non de la mécanique.

Le standard suit l'âge, non la position. Un enfant de 30 mois mesuré couché est malgré tout évalué selon le standard de taille, avec l'ajustement appliqué.

L'ajustement porte sur la mesure, non sur le standard. Vous convertissez la mesure en celle que le standard attend.

L'arrondi est à une décimale, des deux côtés, en un seul endroit. Un flottant issu d'un ajustement n'est pas le même flottant que celui lu dans un CSV, ce qui est la jointure sur clé inexacte à laquelle le cours sur les jointures a consacré un laboratoire.

2.5 Le calcul

```
reference["key"] = list(zip(reference["sex"], reference["lorh"],
                           reference["lenhei"].round(1)))
lms = reference.set_index("key")[["l", "m", "s"]]

def whz(row):
    key = lookup_key(row)
```

```

    if key not in lms.index or pd.isna(row["weight_kg"]):
        return None
    l, m, s = lms.loc[key]
    return (((row["weight_kg"] / m) ** 1) - 1) / (1 * s)

smart["whz"] = smart.apply(whz, axis=1)
print(f"{smart['whz'].notna().sum()} of {len(smart)} computable")

# In R, use the official package rather than the table:
# anthro::anthro_zscores(sex = ..., age = ..., weight = ..., lenhei = ..., measure = ...)

```

864 sur 930 calculables. Les 66 restants se répartissent en trois groupes et la distinction compte.

- **Poids ou âge manquant** — 32 enfants, et ce sont des données manquantes.
- **Taille hors de la plage de référence** — le standard de longueur va de 45,0 à 110,0 cm et celui de taille de 65,0 à 120,0 cm. Un enfant hors de cette plage est réellement inclassable.
- **Une mesure dans la mauvaise unité** — quatre tailles enregistrées en mètres, que le laboratoire du cours sur les jointures a trouvées par anti-jointure. Celles-là sont récupérables.

Rapportez les trois séparément. « 66 enfants exclus » masque le fait que certains étaient des erreurs corrigées.

2.6 Le signalement — deux règles, deux réponses

Les scores z extrêmes sont exclus avant toute prévalence, et il existe deux conventions.

- **Valeurs aberrantes OMS** — bornes fixes, exclure les z hors de -5 à +5. Absolu, comparable entre enquêtes.
- **Valeurs aberrantes SMART** — relatif, exclure les observations à plus de 3 ET de la moyenne *de l'enquête elle-même*. S'adapte à l'enquête, et la réduit davantage quand elle est plus bruitée.

```

who_flagged = smart["whz"].between(-5, 5)

mean, sd = smart["whz"].mean(), smart["whz"].std()
smart_flagged = smart["whz"].between(mean - 3 * sd, mean + 3 * sd)

print(f"WHO flags keep {who_flagged.sum()}, SMART flags keep {smart_flagged.sum()}")

smart <- smart |>
  mutate(who_ok = between(whz, -5, 5),
         smart_ok = between(whz, mean(whz, na.rm = TRUE) - 3 * sd(whz, na.rm = TRUE),
                           mean(whz, na.rm = TRUE) + 3 * sd(whz, na.rm = TRUE)))

```

Elles excluent des enfants différents et donnent des taux légèrement différents. **Dites laquelle vous avez employée** ; un rapport de plausibilité SMART l'exige, et deux enquêtes employant des règles différentes ne sont pas directement comparables.

2.7 L'écart-type est déjà un constat

```
analysable = smart.loc[who_flagged, "whz"]
print(f"n = {len(analysable)}, mean = {analysable.mean():.2f}, "
      f"sd = {analysable.std():.2f}")
```

```
smart |> filter(who_ok) |> summarise(n = n(), mean = mean(whz), sd = sd(whz))
```

852 enfants, **moyenne -0,67, écart-type 1,22**.

La moyenne n'a rien de remarquable — une population un peu sous la médiane de référence, ce qui est normal dans ce secteur. L'écart-type, si. Une enquête bien mesurée produit un écart-type de poids-pour-taille compris entre environ **0,8 et 1,2**, car l'écart-type de la population de référence vaut 1 par construction et les populations réelles ne sont que légèrement plus dispersées.

1,22 se situe à la limite de l'acceptable, et un écart-type au-dessus de la plage signifie l'une de trois choses — une erreur de mesure qui gonfle la dispersion, une population réellement hétérogène, ou un problème de saisie. **C'est le nombre le plus informatif d'un rapport de plausibilité**, et la leçon 5 apprend à le lire correctement.

Remarquez aussi ce qu'il fait à la prévalence. Une distribution plus large place davantage d'enfants au-delà d'un seuil fixe, si bien qu'un écart-type gonflé fait monter la MAG sans qu'aucun enfant ne soit plus malnutri.

2.8 Enregistrez les scores z avec leurs entrées

```
smart[["child_id", "cluster", "team", "age_months", "sex", "weight_kg",
       "height_cm", "measured_lying", "oedema", "whz"]].to_csv(
    "outputs/smart_with_zscores.csv", index=False)
```

```
readr::write_csv(smart, here::here("outputs", "smart_with_zscores.csv"))
```

Gardez les entrées à côté du résultat. Quand quelqu'un conteste une prévalence, la discussion ne porte presque jamais sur l'arithmétique — elle porte sur l'ajustement, la règle de signalement ou les enfants exclus, et les trois ne sont récupérables que si les entrées ont voyagé avec le résultat.

2.9 La suite

Vous avez un score z pour chaque enfant analysable. L'unité suivante en fait une classification — les définitions de cas de la malnutrition aiguë sévère et modérée, et le signe clinique qui prime sur les deux.

Deuxième partie

Qui est un cas

CHAPITRE 3

Les définitions de cas, et la clause que tout le monde oublie

Leçon 3 sur 8 · 120 min

Malnutrition aiguë sévère, modérée et globale, les seuils dont le secteur a convenu, et la clause des œdèmes qui rend fausse toute implémentation en « décompte sous le seuil ».

3.1 Les définitions

Trois termes, et ils s'emboîtent.

Terme	Score z poids-taille	PB	Œdèmes
MAS — malnutrition aiguë sévère	sous -3	sous 115 mm	présents
MAM — malnutrition aiguë modérée	-3 à moins de -2	115 à moins de 125 mm	—
MAG — malnutrition aiguë globale	sous -2	sous 125 mm	présents

La MAG est la MAS plus la MAM. Ce n'est pas une affection distincte, c'est la charge totale de malnutrition aiguë, et c'est le chiffre auquel se réfèrent les phases IPC et le seuil d'urgence de 15 %.

Deux points structurels avant tout code.

Chaque ligne est à elle seule une définition complète. Un enfant est en MAS par le poids-pour-taille, *ou* par le PB, *ou* par les œdèmes — les critères sont des alternatives, non des exigences cumulatives. La leçon 4 porte entièrement sur l'ampleur du désaccord entre les deux premiers.

La colonne des œdèmes est un OU, non un ET. Un enfant présentant des œdèmes bilatéraux prenant le godet est sévère quoi que disent ses mesures. Cette clause est un terme d'une expression booléenne, et c'est le terme le plus souvent omis.

3.2 L'implémentation fausse

```
# WRONG: the definition everybody writes first
sam_wrong = (smart["whz"] < -3).sum()
```

```
# WRONG
sum(smart$whz < -3, na.rm = TRUE)
```

Elle est fausse de deux façons à la fois. Elle écarte les enfants œdémateux, et le `na.rm` ou une comparaison propageant les valeurs nulles écarte silencieusement les enfants sans score z — si bien que le numérateur et le dénominateur diffèrent tous deux de ce que le libellé prétend.

3.3 L'implémentation juste

```
import pandas as pd

analysable = smart[smart["whz"].between(-5, 5)].copy()
oedema = analysable["oedema"] == True

analysable["sam"] = (analysable["whz"] < -3) | oedema
analysable["mam"] = (analysable["whz"].between(-3, -2, inclusive="left")) & ~oedema
analysable["gam"] = analysable["sam"] | analysable["mam"]

n = len(analysable)
for label in ["sam", "mam", "gam"]:
    print(f"{label.upper():4} {analysable[label].sum():>4} / {n}  "
          f"{analysable[label].mean():.1%}")

analysable <- smart |> filter(between(whz, -5, 5))

analysable <- analysable |>
  mutate(sam = whz < -3 | oedema,
         mam = whz >= -3 & whz < -2 & !oedema,
         gam = sam | mam)

analysable |> summarise(across(c(sam, mam, gam), list(n = sum, rate = mean)), n = n())
```

852 enfants analysables : **MAG 14,9 %**, **MAS 3,8 %**.

Trois détails de ce code portent le résultat.

La MAM exclut explicitement les œdèmes. Un enfant œdémateux dont le score z tombe dans la bande modérée est sévère et non modéré, et sans le `& ~oedema` il serait compté dans les deux — si bien que MAS plus MAM dépasserait la MAG.

Les bornes sont semi-ouvertes. $-3 \leq z < -2$. Un enfant exactement à $-3,0$ est sévère. Se tromper ici déplace une poignée d'enfants et c'est le genre de point sur lequel deux analystes découvrent leur désaccord au pire moment.

Le dénominateur est énoncé. 852, non 930 — la différence est faite des enfants sans score z calculable et de ceux hors des bornes de signalement, et la leçon 2 exigeait de rapporter ces trois groupes séparément.

3.4 Vérifiez que les parties totalisent le tout

```
assert (analysable["sam"] & analysable["mam"]).sum() == 0, "a child is both"
assert analysable["gam"].sum() == analysable["sam"].sum() + analysable["mam"].sum()

stopifnot(sum(analysable$sam & analysable$mam) == 0,
          sum(analysable$gam) == sum(analysable$sam) + sum(analysable$mam))
```

Deux assertions, et elles attrapent l'erreur des œdèmes, celle des bornes et toute modification future qui casserait l'emboîtement. Exécutez-les à chaque fois.

3.5 Les définitions par le PB, sur le registre

Le registre de dépistage a le PB et les œdèmes, sans poids ni taille, si bien qu'il ne soutient que les définitions fondées sur le PB.

```
muac = pd.read_csv("muac-screening-artibonite-2024.v1.csv")
measured = muac[muac["muac_mm"].notna()].copy()
oed = measured["oedema"] == True

measured["sam"] = (measured["muac_mm"] < 115) | oed
measured["gam"] = (measured["muac_mm"] < 125) | oed

print(f"n = {len(measured):,}  GAM {measured['gam'].mean():.1%}  "
      f"SAM {measured['sam'].mean():.1%}")

muac |>
  filter(!is.na(muac_mm)) |>
  summarise(n = n(),
            gam = mean(muac_mm < 125 | oedema),
            sam = mean(muac_mm < 115 | oedema))
```

Nommez la mesure dans l'indicateur. `gam_muac_percent` et `gam_whz_percent` sont deux indicateurs différents que l'on appelle tous deux « MAG », et la règle du cours sur la conception d'indicateurs consistant à mettre la mesure dans le nom n'est pas ici du pédantisme — la leçon suivante montre les deux en désaccord sur les trois quarts des cas sévères.

3.6 Ce qu'une définition ne dit pas

Trois choses que la définition de cas laisse délibérément de côté, et que quelqu'un supposera.

L'âge. Les définitions s'appliquent aux enfants de 6 à 59 mois. En dessous de six mois, les normes et les définitions diffèrent ; au-dessus de 59 mois, elles ne s'appliquent pas. Vérifiez la plage d'âge avant de classer.

```
out_of_scope = ~analysable["age_months"].between(6, 59)
print(f"{out_of_scope.sum()} children outside 6-59 months")

sum(!between(analysable$age_months, 6, 59), na.rm = TRUE)
```

L'admission. Une définition de cas dit qui *est* malnutri. Savoir s'il est admis dépend du protocole du programme, qui peut employer un critère, les deux, ou le PB pour le dépistage communautaire et le poids-pour-taille au site. Le champ `admission_criterion` du registre enregistre ce que le site a saisi, non tout ce que l'enfant remplissait.

La gravité à l'intérieur du sévère. Une MAS avec complications exige des soins hospitaliers ; sans complications, ambulatoires. La définition ne les distingue pas, et le test de l'appétit et les signes cliniques qui le font ne figurent pas dans un jeu de données anthropométrique.

Une définition de cas est une règle de classification, non une décision thérapeutique. Confondre les deux produit une charge de cas qui ne correspond aux admissions d'aucun programme et une discussion que les données ne peuvent pas trancher.

3.7 La suite

Vous savez classer un enfant par l'une ou l'autre mesure. La leçon suivante place les deux mesures sur les mêmes enfants et constate qu'elles divergent sur les trois quarts des cas sévères — et que ce désaccord a une cause qui tient en un nombre.

CHAPITRE 4

Là où les deux critères divergent

Leçon 4 sur 8 · 150 min

468 enfants sévères selon le seul poids-pour-taille, 24 selon le seul PB, 142 selon les deux. Une concordance de 22,4 %, et tout le désaccord tient à l'âge — 31,4 mois contre 14,3.

4.1 La question à laquelle le registre peut répondre

Les deux critères d'admission figurent dans les définitions de cas et les deux s'emploient sur le terrain. Savoir s'ils trouvent les mêmes enfants est une question empirique, et elle exige un registre portant les deux mesures sur le même enfant.

Le registre d'admissions PCIMA a exactement cela — PB et poids-pour-taille à l'admission, pour 1 100 épisodes. Cette leçon est ce qu'il dit.

4.2 Calculez les deux sur les mêmes enfants

```
import pandas as pd

cmam = pd.read_csv("cmam-admissions-2024.v1.csv")
cmam["whz"] = cmam.apply(whz, axis=1)          # from lesson 2

both = cmam[cmam["whz"].notna()].copy()
print(f"{len(both)} of {len(cmam)} admissions have both measurements")

both["sam_whz"] = both["whz"] < -3
both["sam_muac"] = both["muac_admission_mm"] < 115

both <- cmam |> filter(!is.na(whz))

both <- both |>
  mutate(sam_whz = whz < -3,
         sam_muac = muac_admission_mm < 115)
```

1 028 sur 1 100. Les 72 sans taille sont le défaut de taille manquante des défauts connus du jeu de données, et c'est le premier constat — **les deux critères n'ont pas le même dénominateur, même sur un registre unique**, car l'un exige deux mesures et l'autre une seule.

4.3 Le tableau croisé

```
table = pd.crosstab(both["sam_muac"], both["sam_whz"],
                    rownames=["SAM by MUAC"], colnames=["SAM by WHZ"])
print(table)

both |> count(sam_muac, sam_whz)
```

	Non sévère par WHZ	Sévère par WHZ
Non sévère par PB	394	468
Sévère par PB	24	142

Lisez les deux cellules en gras. **468 enfants sont sévères par le poids-pour-taille et pas par le PB. 24 le sont par le PB et pas par le poids-pour-taille.** 142 le sont par les deux.

```
overlap = table.loc[True, True]
union = table.values.sum() - table.loc[False, False]
print(f"concordance {overlap / union:.1%} of {union} severe by either measure")
```

```
both |>
  summarise(overlap = sum(sam_muac & sam_whz),
            union = sum(sam_muac | sam_whz)) |>
  mutate(concordance = overlap / union)
```

22,4 %. Des enfants que ce registre qualifierait de sévères selon l'une ou l'autre mesure, moins d'un quart le sont selon les deux.

Ce nombre n'est pas un artefact de ce jeu de données. Les comparaisons publiées trouvent des recouvrements de 10 % à 40 % selon la population, et un programme qui n'a pas mesuré le sien en suppose un.

4.4 Le désaccord, c'est l'âge

```
mean_age = pd.Series({
  "MUAC only": both.loc[both["sam_muac"] & ~both["sam_whz"], "age_months"].mean(),
  "WHZ only": both.loc[both["sam_whz"] & ~both["sam_muac"], "age_months"].mean(),
  "Both": both.loc[both["sam_muac"] & both["sam_whz"], "age_months"].mean(),
})
print(mean_age.round(1))
```

```
both |>
  mutate(group = case_when(sam_muac & sam_whz ~ "both",
                           sam_muac ~ "muac only",
                           sam_whz ~ "whz only",
                           TRUE ~ "neither")) |>
  summarise(mean_age = mean(age_months), n = n(), .by = group)
```

Groupe	Âge moyen
Sévère par le seul PB	14,3 mois
Sévère par les deux	15,6 mois
Sévère par le seul poids-pour-taille	31,4 mois

Voilà. Les enfants que le PB trouve et que le poids-pour-taille manque ont la moitié de l'âge de ceux que le poids-pour-taille trouve et que le PB manque.

Le mécanisme est simple et ce n'est un défaut d'aucune des deux mesures. Le PB croît avec l'âge : un bras en bonne santé mesure environ 135 mm à six mois et environ 155 mm à cinq ans. Un seuil fixe de 115 mm représente donc un déficit bien plus profond pour un enfant de cinq ans que pour un nourrisson, si bien que le PB devient progressivement plus

difficile à franchir à mesure que l'enfant grandit. Le poids-pour-taille n'a aucun terme d'âge et s'applique également sur toute la plage.

```
by_age = both.assign(
    band=pd.cut(both["age_months"], [5, 12, 24, 36, 60],
                labels=["6-11", "12-23", "24-35", "36-59"])
).groupby("band")[["sam_muac", "sam_whz"]].mean()
print((by_age * 100).round(1))

both |>
  mutate(band = cut(age_months, c(5, 12, 24, 36, 60),
                    labels = c("6-11", "12-23", "24-35", "36-59"))) |>
  summarise(across(c(sam_muac, sam_whz), mean), .by = band)
```

Exécutez-le et les deux courbes se croisent. Le PB identifie davantage les plus jeunes ; le poids-pour-taille davantage les plus âgés. **Un programme admettant sur le seul PB fait tourner une charge de cas plus jeune qu'un programme admettant sur le seul poids-pour-taille**, et c'est un fait clinique sur qui est traité, non un détail de mesure.

4.5 Lequel a raison ?

Aucun, et la question est la mauvaise. Ce à quoi chacun *sert* diffère.

- **Le PB prédit la mortalité au moins aussi bien**, et mieux dans plusieurs études. Si le but de l'admission est de traiter les enfants les plus susceptibles de mourir, c'est un argument fort.
- **Le poids-pour-taille est le standard des enquêtes**, et les seuils IPC comme les comparaisons internationales reposent sur lui. Si le but est de classer une population, c'est décisif.
- **Le PB est opérationnellement praticable** au niveau communautaire. Le poids-pour-taille ne l'est pas.

La réponse du secteur est d'employer les deux, d'admettre sur l'un ou l'autre, et de décharger sur celui qui a servi à l'admission. Ce qui est cliniquement sensé et crée le problème analytique pour lequel cette leçon existe.

4.6 Ce que cela implique pour vos chiffres

Quatre conséquences, et chacune est une question qui vous sera posée.

La charge de cas dépend du critère. Admettre sur l'un ou l'autre donne ici une charge de 634 enfants sévères ; sur le seul PB, 166 ; sur le seul poids-pour-taille,

1. Ce ne sont pas des estimations de la même quantité.

La prévalence n'est pas comparable entre mesures. Une MAG par le PB et une MAG par le poids-pour-taille sont deux indicateurs différents. La leçon 6 montre ce que cela produit quand un chiffre est lu face au seuil de 15 %.

La sortie doit employer le critère d'admission. Un enfant admis sur le PB et déchargé sur le poids-pour-taille peut sortir avant guérison, ou être gardé bien après. Le protocole PCIMA est explicite et les registres, régulièrement, ne le sont pas.

Un programme qui change de critère casse sa propre série. Une charge de cas qui monte de 40 % au changement de protocole est une rupture de définition, et la règle du cours sur la conception d'indicateurs à propos du versionnement s'applique exactement.

```
for name, mask in [("either", both["sam_muac"] | both["sam_whz"]),
                  ("muac only", both["sam_muac"]),
                  ("whz only", both["sam_whz"])]:
    print(f"{name:10} caseload {mask.sum():>4}")
```

```
both |> summarise(either = sum(sam_muac | sam_whz),
                 muac = sum(sam_muac), whz = sum(sam_whz))
```

4.7 Rapportez toujours la paire

Severe acute malnutrition, admissions 2024, n = 1,028 with both measurements

By MUAC (<115 mm)	166	
By weight-for-height (z < -3)	610	
By either	634	
By both	142	(22.4% of those severe by either)

Children severe by MUAC alone average 14.3 months; by weight-for-height alone, 31.4 months. The two criteria identify overlapping but substantially different caseloads, and the difference is age.

Six lignes. Elles devancent la question, nomment le mécanisme, et empêchent quelqu'un de comparer une charge fondée sur le PB à une charge fondée sur le poids-pour-taille comme si l'écart était une évolution nutritionnelle.

4.8 La suite

Vous savez classer des enfants et vous savez de quoi la classification dépend. L'unité suivante se demande si l'enquête ayant produit les mesures est croyable — le rapport de plausibilité SMART, où l'écart-type de 1,22 de la leçon 2 est enfin lu.

Troisième partie

Juger une enquête

CHAPITRE 5

Cette enquête est-elle croyable ?

Leçon 5 sur 8 · 150 min

Le rapport de plausibilité SMART, sur une enquête qui échoue à deux de ses contrôles. Un écart-type de 1,22, une équipe mesurant 0,6 score z trop bas, 69 % de ses tailles sur un demi-centimètre, et un quart des âges sur une année entière.

5.1 Une enquête est acceptée ou ne l'est pas

La méthodologie SMART fait quelque chose d'inhabituel et de précieux — elle publie un ensemble de contrôles qu'une enquête doit passer avant que sa prévalence soit employée, et ces contrôles se calculent sur les données de l'enquête elle-même.

Cela signifie qu'une enquête peut être **rejetée sur ses propres preuves**, avant toute discussion sur le sens du nombre. Cette leçon applique les contrôles à l'enquête SMART de la plateforme, qui en passe certains et en échoue d'autres.

5.2 Les contrôles

Contrôle	Ce qu'il détecte	Acceptable
Valeurs signalées	Mesures impossibles ou extrêmes	sous 2,5 % environ
Écart-type du WHZ	Erreur de mesure gonflant la dispersion	0,8 à 1,2
Préférence de chiffres	Arrondi plutôt que lecture	faible, et homogène entre équipes
Attraction des âges	Âge estimé plutôt que documenté	faible
Sex-ratio	Biais de sélection dans qui a été mesuré	proche de 1,0
Biais d'équipe	Une équipe qui mesure autrement	moyennes proches entre équipes

5.3 Les signalements

```
computable = smart["whz"].notna()
flagged = computable & ~smart["whz"].between(-5, 5)

print(f"{computable.sum()} computable, {flagged.sum()} flagged "
      f"({flagged.sum() / computable.sum():.1%})")
```

```
smart |>
  filter(!is.na(whz)) |>
  summarise(n = n(), flagged = sum(!between(whz, -5, 5)),
            rate = flagged / n)
```

12 sur 864, environ 1,4 %. **Passe.** Un taux de signalement au-dessus de quelques pour cent signale des problèmes de mesure ou de saisie assez graves pour mettre en doute tout le reste, si bien que ce contrôle passe en premier et conditionne les autres.

5.4 L'écart-type

```
analysable = smart.loc[smart["whz"].between(-5, 5), "whz"]
print(f"n = {len(analysable)}, mean = {analysable.mean():.2f}, "
      f"sd = {analysable.std():.2f}")
```

```
smart |> filter(between(whz, -5, 5)) |> summarise(n = n(), sd = sd(whz))
```

1,22. Hors de la bande acceptable, à sa limite haute.

C'est le nombre le plus important d'un rapport de plausibilité et il vaut la peine d'être précis sur les raisons. La population de référence de l'OMS a un écart-type de 1 par construction. Une population réelle est légèrement plus dispersée, si bien que 1,0 à 1,2 est ce que produit une enquête bien mesurée. Au-delà, la dispersion est gonflée par quelque chose, et il y a trois candidats.

- **L'erreur de mesure.** Une erreur aléatoire sur le poids ou la taille élargit la distribution sans en déplacer le centre. La cause la plus fréquente.
- **Une population réellement hétérogène.** Deux sous-populations très différentes échantillonnées ensemble.
- **L'erreur de saisie.** Chiffres inversés, mauvaises unités, enregistrements mélangés.

Un écart-type gonflé fait monter la prévalence. Une distribution plus large place davantage d'enfants au-delà d'un seuil fixe, si bien que la MAG monte sans qu'un seul enfant ne soit plus malnutri. C'est la raison d'être du contrôle et la raison pour laquelle il n'est pas facultatif.

5.5 La préférence de chiffres

Une mesure lue sur un instrument a un dernier chiffre quasi uniforme. Une mesure arrondie à l'œil non.

```
smart["terminal"] = (smart["height_cm"] * 10).round() % 10
by_team = (
    smart.assign(rounded=smart["terminal"].isin([0, 5]))
    .groupby("team")
    .agg(rounded=("rounded", "mean"), n=("rounded", "size"))
)
print((by_team["rounded"] * 100).round(0))
```

```
smart |>
  mutate(rounded = round(height_cm * 10) %% 10 %in% c(0, 5)) |>
  summarise(share = mean(rounded), n = n(), .by = team)
```

Équipe	Tailles finissant par ,0 ou ,5
1	18 %
2	69 %
3	18 %
4	23 %

Deux chiffres sur dix font 20 % en l'absence de préférence. Les équipes 1, 3 et 4 y sont. **L'équipe 2 est à 69 %**, trois fois et demie les autres.

Échec, et cela n'exige aucun test de signification — la discipline d'étalonnage du cours d'évaluation s'applique aux résultats marginaux, et celui-ci ne l'est pas. Le mécanisme se nomme : l'équipe a lu la toise au demi-centimètre le plus proche au lieu du millimètre.

Notez ce que cela produit. Arrondir au 0,5 cm ajoute du bruit à la taille, qui se propage au poids-pour-taille, qui gonfle l'écart-type — si bien que les deux contrôles échoués sont liés et que l'un cause en partie l'autre.

5.6 L'attraction des âges

```
ages = smart["age_months"].dropna()
whole_years = (ages % 12 == 0).mean()

print(f"{whole_years:.0%} of ages on an exact whole year")
print(ages.value_counts().reindex([22, 23, 24, 25, 26]).to_dict())
```

```
smart |>
  filter(!is.na(age_months)) |>
  summarise(whole_years = mean(age_months %% 12 == 0))
```

24 % sur une année entière, contre environ 9 % attendus. Soixante-six enfants à exactement 24 mois, contre 15 et 19 de part et d'autre.

Échec, et la cause n'est presque jamais la négligence — c'est que les actes de naissance ne sont pas universels et qu'un accompagnant interrogé sur l'âge d'un enfant répond en années. L'action corrective est un calendrier d'événements locaux et un questionnement croisé, non une réprimande.

Cela compte ici pour une raison précise. Les normes de l'OMS changent de table de référence à 24 mois et les groupes d'âge SMART sont bornés aux multiples de 12 mois, si bien que l'attraction dépose un paquet d'enfants exactement là où la classification change.

5.7 Sex-ratio et biais d'équipe

```
ratio = (smart["sex"] == "m").sum() / (smart["sex"] == "f").sum()
print(f"sex ratio m:f = {ratio:.2f}")

by_team = (
  smart[smart["whz"].between(-5, 5)]
  .groupby("team")
  .agg(mean_whz=("whz", "mean"), n=("whz", "size"))
)
print(by_team.round(2))
```

```
smart |>
  filter(between(whz, -5, 5)) |>
  summarise(mean_whz = mean(whz), n = n(), .by = team)
```

Équipe	WHZ moyen	n
1	-0,69	204
2	-0,42	228
3	-1,09	224
4	-0,44	196

Sex-ratio proche de 1,0 — **passe**.

Les moyennes d'équipe s'étalent de -0,42 à -1,09. **Échec**. Les grappes ont été attribuées aux équipes indépendamment de l'état nutritionnel, si bien qu'un écart de 0,67 score z entre équipes n'est pas une différence entre les enfants, c'est une différence de mesure.

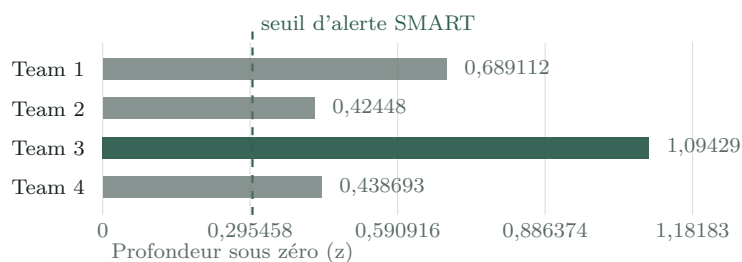


FIGURE 5.1 – Score z poids-pour-taille moyen par équipe de mesure, tracé en profondeur sous zéro. Les grappes ayant été attribuées aux équipes indépendamment de l'état nutritionnel, un écart de cette ampleur entre équipes est un défaut de mesure et non une différence réelle entre populations.

Le coût est direct.

```
gam_by_team = (
    smart[smart["whz"].between(-5, 5)]
    .assign(gam=lambda d: (d["whz"] < -2) | (d["oedema"] == True))
    .groupby("team")["gam"].mean()
)
print((gam_by_team * 100).round(1))
```

```
smart |>
  filter(between(whz, -5, 5)) |>
  summarise(gam = mean(whz < -2 | oedema), .by = team)
```

Les grappes de l'équipe 3 donnent une MAG de 22,3 % contre 9,2 % à 16,2 % pour les autres. **L'équipe 3 apporte un quart de l'échantillon et tire l'estimation d'ensemble vers le haut d'environ deux points.**

5.8 Le verdict

```
plausibility = pd.DataFrame({
    "check": ["Flagged records", "SD of WHZ", "Digit preference",
              "Age heaping", "Sex ratio", "Team bias"],
    "value": ["1.4%", "1.22", "69% (team 2)", "24% whole years", "1.05",
              "-0.42 to -1.09"],
    "acceptable": ["<2.5%", "0.8-1.2", "even across teams", "low", "~1.0",
                  "close across teams"],
    "verdict": ["pass", "fail", "fail", "fail", "pass", "fail"],
})
```

```
| print(plausibility)

tibble::tribble(
  ~check,          ~value,          ~verdict,
  "Flagged records", "1.4%",          "pass",
  "SD of WHZ",      "1.22",          "fail",
  "Digit preference", "69% (team 2)",  "fail",
  "Age heaping",    "24%",          "fail",
  "Sex ratio",      "1.05",          "pass",
  "Team bias",      "-0.42 to -1.09", "fail"
)
```

Quatre échecs. **Ce qui suit est un jugement, non un résultat arithmétique**, et les options honnêtes sont au nombre de trois.

- **Rejeter l'enquête.** Défendable avec quatre échecs, et coûteux — six semaines et un budget, perdus.
- **Accepter avec des limites énoncées.** Rapporter la prévalence avec le tableau de plausibilité à côté et une mention explicite que l'écart-type et le biais d'équipe poussent tous deux l'estimation vers le haut. Le plus fréquent, et cela exige que le tableau voyage avec le chiffre.
- **Réanalyser en excluant l'équipe 3.** Défendable seulement si vous le dites bien en vue, et cela coûte un quart de l'échantillon et élargit donc l'intervalle.

Ce qu'il ne faut pas faire est de publier la prévalence sans le rapport de plausibilité. Le nombre n'est pas faux ; il est sans qualification, et la qualification change ce qu'il soutient.

5.9 Rédigez sans accuser personne

N'écrivez pas	Écrivez
« Les mesures de l'équipe 2 ne sont pas fiables »	« Les tailles de l'équipe 2 se terminent par ,0 ou ,5 dans 69 % des
« L'équipe 3 a falsifié des mesures »	« Le poids-pour-taille moyen de l'équipe 3 est de 0,4 à 0,7 score z
« Les âges sont inexacts »	« 24 % des âges tombent sur une année entière exacte contre 9 %

La colonne de droite nomme ce qui a été observé, son ampleur et un mécanisme. C'est aussi ce qui rend l'action corrective évidente — réétalonner une balance, reformer à la lecture de la toise, ajouter un calendrier d'événements — là où la colonne de gauche produit une équipe sur la défensive et de moins bonnes données au tour suivant.

5.10 La suite

Vous savez si l'enquête est croyable et sous quelles réserves. La leçon suivante en tire le nombre pour lequel elle existe — prévalences de la MAG et de la MAS, avec un intervalle tenant compte du plan en grappes, lues au regard des phases IPC.

CHAPITRE 6

L'intervalle qui chevauche deux phases IPC

Leçon 6 sur 8 · 150 min

La MAG vaut 14,9 %, au sommet de la phase 3 de l'IPC. L'intervalle ajusté du plan va de 11,1 % à 18,7 %, ce qui couvre les phases 3 et 4 — et l'enquête ne peut pas dire dans laquelle se trouve la population.

6.1 Le nombre pour lequel l'enquête existe

Tout ce qui précède était préparatoire. L'enquête a été financée pour répondre à une question — quelle est l'ampleur de la malnutrition aiguë — et cette leçon produit la réponse sous la forme qu'elle doit prendre : une prévalence, un intervalle et une phase.

6.2 L'estimation, avec le plan

Une enquête SMART est une enquête en grappes. Le cours sur les enquêtes a établi ce que cela implique ; ici cela arrive avec un seuil attaché.

```
import numpy as np
import pandas as pd

analysable = smart[smart["whz"].between(-5, 5)].copy()
analysable["gam"] = (analysable["whz"] < -2) | (analysable["oedema"] == True)
analysable["sam"] = (analysable["whz"] < -3) | (analysable["oedema"] == True)

def cluster_prevalence(df, column, cluster="cluster"):
    n = len(df)
    p = df[column].mean()

    totals = df.groupby(cluster)[column].agg(["sum", "size"])
    residual = totals["sum"] - p * totals["size"]
    m = len(totals)

    variance = m / (m - 1) * (residual**2).sum() / n**2
    se = np.sqrt(variance)
    srs = p * (1 - p) / n
    return {"p": p, "se": se, "deff": variance / srs, "clusters": m, "n": n}

for column in ["gam", "sam"]:
    r = cluster_prevalence(analysable, column)
    t = 2.045 # t(0.975, 29 df)
    print(f"{column.upper()}: {r['p']:.1%} 95% CI {r['p'] - t * r['se']:.1%} to "
          f"{r['p'] + t * r['se']:.1%} deff {r['deff']:.2f} n {r['n']}")

library(survey)

design <- svydesign(ids = ~cluster, weights = NULL, data = analysable)

svyciprop(~I(whz < -2 | oedema), design, method = "logit")
svymean(~I(whz < -2 | oedema), design, deff = TRUE)
```

	Estimation	IC 95 %	Effet de plan	n
MAG	14,9 %	11,1 – 18,7 %	2,29	852
MAS	3,8 %	2,3 – 5,2 %	1,20	852

Deux points à retenir de ce tableau avant de lire la phase.

L'effet de plan vaut 2,29 pour la MAG et 1,20 pour la MAS. Deux indicateurs, une enquête, deux effets de plan — car la malnutrition se regroupe géographiquement et la malnutrition sévère, plus rare, se regroupe de façon moins détectable. Les 852 enfants de l'enquête en valent environ 372 indépendants pour la MAG.

L'intervalle naïf serait de 12,5 à 17,3 %. Ignorer les grappes rétrécit l'intervalle de 37 %, et la section suivante dit ce que ce rétrécissement aurait coûté.

6.3 Les phases IPC

L'IPC classe la prévalence de malnutrition aiguë en cinq phases, et ce sont elles qui transforment un pourcentage en décision.

Phase	MAG par poids-pour-taille
1 — Acceptable	sous 5 %
2 — Alerte	5 à 9,9 %
3 — Sérieuse	10 à 14,9 %
4 — Critique	15 à 29,9 %
5 — Extrêmement critique	30 % et plus

```
def ipc_phase(gam):
    for threshold, phase in [(0.05, 1), (0.10, 2), (0.15, 3), (0.30, 4)]:
        if gam < threshold:
            return phase
    return 5

r = cluster_prevalence(analysable, "gam")
t = 2.045
low, high = r["p"] - t * r["se"], r["p"] + t * r["se"]

print(f"point estimate: phase {ipc_phase(r['p'])}")
print(f"interval spans: phase {ipc_phase(low)} to phase {ipc_phase(high)}")

c(point = 0.149, low = 0.111, high = 0.187)
```

L'estimation ponctuelle vaut 14,9 %, soit la phase 3 — à un dixième de point près.

L'intervalle va de 11,1 % à 18,7 %, ce qui couvre les phases 3 et 4.

6.4 Ce que cela veut dire, dit franchement

C'est le moment vers lequel tout le cours tendait, et l'énoncé honnête est inconfortable.

La malnutrition aiguë globale est estimée à 14,9 % (IC 95 % 11,1-18,7). L'estimation ponctuelle relève de la phase 3 de l'IPC (Sérieuse). L'intervalle de confiance couvre les phases 3 et 4 (Critique), si bien que **cette enquête ne permet pas de déterminer**

dans quelle phase se trouve la population.

Trois choses que fait cette phrase.

Elle refuse d'arrondir jusqu'au seuil. 14,9 % rapporté « environ 15 % » a pris une décision de classification par arrondi, et la différence entre phase 3 et phase 4 est une différence de réponse.

Elle énonce l'intervalle avant la phase. Un lecteur qui voit « phase 3 » d'abord ne la révisera pas en arrivant à l'intervalle.

Elle dit que l'enquête ne peut pas trancher. C'est le verdict à trois branches sur lequel insistait le cours sur les enquêtes, et ici la troisième branche est la bonne.

6.5 Ce qui se serait passé sans l'effet de plan

```
srs_se = np.sqrt(r["p"] * (1 - r["p"]) / r["n"])
print(f"naive interval: {r['p'] - 1.96 * srs_se:.1%} to {r['p'] + 1.96 * srs_se:.1%}")
```

```
p <- 0.149; n <- 852
c(p - 1.96 * sqrt(p * (1 - p) / n), p + 1.96 * sqrt(p * (1 - p) / n))
```

12,5 % à 17,3 %. **Toujours à cheval sur les deux phases** — donc sur cette enquête l'effet de plan ne change pas le verdict, et le dire vaut autant qu'un constat.

Mais voyez à quel point c'est passé près. Si l'estimation avait valu 13,5 %, l'intervalle naïf serait resté dans la phase 3 et l'intervalle ajusté du plan serait entré en phase 4. **La mise en grappes décide de la classification dès que l'estimation est à environ deux points d'un seuil**, ce qui, dans ce secteur, est le cas la plupart du temps.

6.6 Lisez-le avec le rapport de plausibilité

La leçon précédente a trouvé quatre contrôles échoués, et deux d'entre eux poussent dans un sens connu.

- **L'écart-type gonflé (1,22) fait monter la MAG.** Une distribution plus large place davantage d'enfants au-delà d'un seuil fixe.
- **Les mesures basses de l'équipe 3 font monter la MAG.** Ses grappes donnent 22,3 % contre 9,2 à 16,2 % pour les autres, et elle apporte un quart de l'échantillon.

L'estimation a donc plus de chances d'être trop haute que trop basse, ce qui pousse la valeur vraie vers la moitié basse de l'intervalle — vers la phase 3 plutôt que la phase 4.

```
without_team3 = analysable[analysable["team"] != 3]
r3 = cluster_prevalence(without_team3, "gam")
print(f"excluding team 3: {r3['p']:.1%} on n={r3['n']}, "
      f"{r3['clusters']} clusters")
```

```
analysable |> filter(team != 3) |> summarise(gam = mean(gam), n = n())
```

Rapportez la sensibilité, ne la substituez pas. Le chiffre principal reste l'estimation sur échantillon complet ; l'exclusion figure à côté comme indice sur le sens. Écarter en silence

un quart d'une enquête parce qu'il donne une réponse gênante est précisément ce que tout ce module existe pour prévenir.

6.7 Le bloc de rapportage complet

```
Global acute malnutrition (weight-for-height  $z < -2$  or oedema)
14.9% 95% CI 11.1-18.7 n = 852 design effect 2.29 30 clusters
IPC Phase 3 by point estimate; interval spans Phase 3 and Phase 4.

Severe acute malnutrition ( $z < -3$  or oedema)
3.8% 95% CI 2.3-5.2 n = 852 design effect 1.20

Plausibility: 4 of 6 SMART checks failed (SD 1.22, digit preference in one team,
age heaping 24%, team bias -0.42 to -1.09). Two of the four push the estimate
upward. Excluding the affected team gives 12.6% on 628 children.

Analysable denominator: 852 of 930 measured. 66 had no computable z-score
(32 missing weight or age, 30 outside the reference range, 4 heights in metres,
recoverable) and 12 were flagged by the WHO bounds.
```

Douze lignes. Elles portent l'estimation, son intervalle, son plan, sa classification, ses biais connus et son dénominateur, et il ne reste rien qu'un relecteur puisse demander sans que le bloc y réponde.

6.8 La suite

L'enquête dit quelle est la situation. L'unité suivante demande ce que le programme en a fait — la performance PCIMA au regard des standards Sphere, et le dénominateur qui déplace le taux de guérison de neuf points.

Quatrième partie

Juger un programme

CHAPITRE 7

Le dénominateur qui déplace le taux de guérison de neuf points

Leçon 7 sur 8 · 150 min

Guéris, abandons, décès, non-réponse — au regard des standards minimums Sphere, sur le dénominateur que Sphere spécifie réellement. 82,3 % ou 73,6 %, selon un choix que personne ne documente, et un site hors standard que le chiffre du district masque.

7.1 Quatre issues et un standard

Un programme PCIMA est jugé sur ce qu'il est advenu des enfants qu'il a admis, et les standards minimums Sphere donnent les seuils.

Issue	Minimum Sphere, soins thérapeutiques ambulatoires
Guéris	plus de 75 %
Décès	moins de 10 %
Abandons	moins de 15 %
Non-réponse	pas de seuil fixe ; investiguée si élevée

L'alimentation supplémentaire a son propre jeu — guéris au-dessus de 75 %, décès sous 3 %, abandons sous 15 %. Les nombres diffèrent ; la structure non.

Chacun est une proportion, et la discussion porte toujours sur le dénominateur.

7.2 Les trois dénominateurs candidats

```
import pandas as pd

cmam = pd.read_csv("cmam-admissions-2024.v1.csv")

print(f"admissions:           {len(cmam):>5}")
print(f"with an outcome:       {cmam['outcome'].notna().sum():>5}")
print(f"excluding transfers:      "
      f"{(cmam['outcome'].notna() & (cmam['outcome'] != 'transferred')).sum():>5}")

cmam |> summarise(
  admissions = n(),
  with_outcome = sum(!is.na(outcome)),
  excluding_transfers = sum(!is.na(outcome) & outcome != "transferred")
)
```

Dénominateur	n	Ce qu'il fait des exclus
Toutes les admissions	1 100	Les enfants encore en traitement comptés comme non guéris
Ayant atteint une issue	1 029	Les transferts comptés comme une issue
Ayant atteint une issue, transferts exclus	984	—

```
cured = (cmam["outcome"] == "cured").sum()
for label, n in [("all admissions", 1100), ("with an outcome", 1029),
                 ("excluding transfers", 984)]:
    print(f"cure rate on {label:22} {cured / n:.1%}")
```

```
cured <- sum(cmam$outcome == "cured", na.rm = TRUE)
c(all = cured / 1100, outcome = cured / 1029, sphere = cured / 984)
```

73,6 %, 78,7 %, 82,3 %. Un numérateur, trois dénominateurs, neuf points.

Et l'un d'eux franchit un seuil : sur toutes les admissions le programme est *sous* le minimum Sphere de 75 % et sur le dénominateur Sphere il est confortablement au-dessus. La différence n'est pas la performance. C'est un dénominateur que personne n'a écrit.

7.3 Celui que Sphere désigne

Les enfants ayant atteint une issue thérapeutique, transferts exclus. 984 ici.

Les deux exclusions ont des motifs différents et les deux sont justes.

Les enfants encore en traitement ne sont pas une issue. Soixante et onze enfants ont été admis assez tard pour ne pas être sortis à la date d'arrêt. Ils ne sont ni guéris ni abandons ; ils sont en traitement. Les compter dans une catégorie quelconque est une affirmation sur un avenir qui n'a pas eu lieu.

Les transferts sont l'issue de quelqu'un d'autre. Quarante-cinq enfants sont passés d'un programme à l'autre — ambulatoire vers hospitalier, ou l'inverse — et leur issue thérapeutique sera enregistrée là où ils aboutissent. Les compter ici reviendrait à les compter deux fois entre les deux programmes, ou à créditer ce programme d'un résultat qu'il n'a pas produit.

```
performance = cmam[cmam["outcome"].notna() & (cmam["outcome"] != "transferred")]

rates = performance["outcome"].value_counts(normalize=True)
print((rates * 100).round(1))
print(f"n = {len(performance)}")
```

```
performance <- cmam |> filter(!is.na(outcome), outcome != "transferred")
```

```
performance |> count(outcome) |> mutate(rate = n / sum(n))
```

Issue	Taux	Sphere	
Guéris	82,3 %	plus de 75 %	passe
Abandons	13,7 %	moins de 15 %	passe
Décès	1,0 %	moins de 10 %	passe
Non-réponse	2,9 %	—	

Quatre lignes, toutes dans le standard. **C'est là que s'arrêtent la plupart des rapports PCIMA, et c'est là que le constat se cache.**

7.4 Décomposez par site

```
by_site = (
    performance.assign(**{o: performance["outcome"] == o
```

```

        for o in ["cured", "defaulted", "died", "non-response"]})
.groupby("site_id")
.agg(n=("outcome", "size"), cured=("cured", "mean"),
     defaulted=("defaulted", "mean"), died=("died", "mean"))
)
print((by_site[["cured", "defaulted", "died"]] * 100).round(1))

```

```

performance |>
  summarise(n = n(),
            cured = mean(outcome == "cured"),
            defaulted = mean(outcome == "defaulted"),
            died = mean(outcome == "died"),
            .by = site_id) |>
  arrange(desc(defaulted))

```

Site	n	Abandons
SITE-03	156	20,5 %
SITE-02	168	15,5 %
SITE-06	125	13,6 %
SITE-05	142	11,3 %
SITE-01	188	11,2 %
SITE-04	205	11,2 %

Le programme abandonne à 13,7 % et un site à 20,5 %. Le chiffre du district est sous le maximum Sphere et ce site ne l'est pas, de cinq points.

C'est toute la raison pour laquelle le cours d'évaluation insistait sur les distributions plutôt que sur les moyennes, et cela arrive ici avec une conséquence clinique — un enfant qui abandonne est un enfant qui a cessé le traitement avant guérison.

```

flagged = by_site[(by_site["defaulted"] > 0.15) | (by_site["cured"] < 0.75)]
print(f"{len(flagged)} of {len(by_site)} sites outside a Sphere standard")

```

```

performance |>
  summarise(cured = mean(outcome == "cured"),
            defaulted = mean(outcome == "defaulted"), .by = site_id) |>
  filter(defaulted > 0.15 | cured < 0.75)

```

Deux sites, et le second est limite à 15,5 %. **Signalez par rapport au standard, non par rapport à la moyenne du district** — c'est au standard que le programme est tenu.

7.5 La cause est en général la distance, non l'observance

L'abandon a une littérature et elle est constante : le meilleur prédicteur est le temps de trajet jusqu'au site. Un accompagnant qui doit perdre une journée de travail chaque semaine pour marcher deux heures cessera de venir quand l'enfant aura meilleure mine, ce qui est rationnel et n'est pas de l'inobservance.

```

length_of_stay = (
  pd.to_datetime(performance["discharge_date"])
  - pd.to_datetime(performance["admission_date"])
).dt.days

print(performance.assign(stay=length_of_stay))

```

```
.groupby("outcome")["stay"].median().round(0))
```

```
performance |>
  mutate(stay = as.integer(as.Date(discharge_date) - as.Date(admission_date))) |>
  summarise(median_stay = median(stay), .by = outcome)
```

Les abandons partent bien plus tôt que les enfants guéris — le séjour médian est d'environ trois semaines contre environ huit. **Un abandon n'est pas un échec thérapeutique, c'est un enfant parti avant la fin du traitement**, et l'action corrective est la décentralisation ou un appui au transport, non du counseling.

C'est la discipline de cause racine du cours d'évaluation, appliquée à un indicateur clinique.

7.6 Le gain de poids, et les saisies qui ne peuvent pas être justes

```
gain = (
  (performance["weight_discharge_kg"] - performance["weight_admission_kg"])
  / performance["weight_admission_kg"]
  / length_of_stay * 1000
)
print(gain.describe().round(1))
print(f"{{gain < 0}.sum()}} discharges with weight loss")
```

```
performance |>
  mutate(gain = (weight_discharge_kg - weight_admission_kg) /
    weight_admission_kg / as.integer(as.Date(discharge_date) -
    as.Date(admission_date)) * 1000) |>
  summarise(median = median(gain, na.rm = TRUE), negative = sum(gain < 0, na.rm = TRUE))
```

Le gain de poids en grammes par kilogramme et par jour est la mesure classique de la réponse au traitement, et neuf sorties présentent une perte de poids. C'est possible chez un enfant décédé ou ayant abandonné tôt, et c'est aussi exactement l'allure d'une saisie inversée. **Le registre ne permet pas de les distinguer**, ce qui est un constat pour la section qualité des données plutôt qu'un nombre à nettoyer.

7.7 Le tableau de performance à publier

CMAM performance, 2024

Admissions	1,100		
Still in treatment at cut-off	71	excluded: not an outcome	
Transferred	45	excluded: outcome recorded elsewhere	
Performance denominator	984		
Cured	82.3%	Sphere >75%	pass
Defaulted	13.7%	Sphere <15%	pass
Died	1.0%	Sphere <10%	pass
Non-response	2.9%		

One site (SITE-03, n=156) defaults at 20.5%, outside the Sphere maximum.
 Median stay for defaulters is 24 days against 58 for cured children,
 consistent with distance rather than non-response to treatment.

Les lignes exclues sont montrées, non supprimées. Un lecteur peut reconstituer n'importe lequel des trois dénominateurs à partir de ce bloc, ce qui fait la différence entre un tableau de performance et une affirmation.

7.8 La suite

Le programme guérit les enfants qu'il admet. La dernière leçon pose la question plus difficile — quelle part des enfants ayant besoin d'un traitement il a jamais vue — et pourquoi le chiffre que la plupart des programmes rapportent comme couverture n'en est pas une.

CHAPITRE 8

Pourquoi admissions sur charge attendue n'est pas la couverture

Leçon 8 sur 8 · 120 min

Le chiffre que la plupart des programmes rapportent comme couverture est le rapport de deux choses dont aucune n'est le nombre d'enfants malnutris. Ce qu'exige réellement la couverture, et quoi rapporter quand on ne peut pas l'estimer.

8.1 Le chiffre que tout le monde rapporte

Quelque part dans presque tout rapport PCIMA figure une ligne de ce genre.

Couverture du programme : 68 %

et en dessous, si l'on cherche, le calcul.

admissions during the period / expected caseload for the period

La charge attendue est elle-même calculée, en général ainsi.

population under five x GAM prevalence x incidence correction x coverage target

```
UNDER_FIVE = 21500
GAM_PREVALENCE = 0.149          # from the SMART survey
INCIDENCE_CORRECTION = 1.6      # new cases arising over the year
COVERAGE_TARGET = 0.50

expected = UNDER_FIVE * GAM_PREVALENCE * INCIDENCE_CORRECTION * COVERAGE_TARGET
admissions = 1100

print(f"expected caseload {expected:.0f}; admissions {admissions:}; "
      f"'coverage' {admissions / expected:.0%}")
```

```
expected <- 21500 * 0.149 * 1.6 * 0.50
c(expected = expected, ratio = 1100 / expected)
```

Ce nombre n'est pas la couverture. Il vaut la peine d'être précis sur ce qu'il est, car il n'est pas inutile — il ne correspond simplement pas à son étiquette.

8.2 Ce qu'il mesure réellement

La couverture est une proportion — enfants traités sur enfants ayant besoin de traitement. La formule ci-dessus n'a ni l'un ni l'autre.

Le numérateur, ce sont des admissions, non des enfants. Un enfant réadmis après rechute compte pour deux admissions. La leçon sur la portée du cours sur les indicateurs faisait ce constat en général; ici il gonfle le numérateur du taux de rechute.

Le dénominateur est une prévision, non un décompte. Chacun de ses termes est une estimation — la population d'une projection de recensement, la prévalence d'une enquête avec son propre intervalle, la correction d'incidence d'une valeur de la littérature, et la cible de couverture d'une proposition. Multipliez quatre estimations et le produit porte les quatre incertitudes.

La cible de couverture est au dénominateur. Regardez-la à nouveau. Un programme visant 50 % de couverture divise par la moitié de la charge, si bien qu'atteindre la cible se lit 100 %. La mesure est construite pour atteindre 100 % quand le plan est tenu, ce qui en fait un **taux de réalisation du plan** — un indicateur de gestion parfaitement raisonnable, sous un nom qui ne lui convient pas.

```
without_target = admissions / (UNDER_FIVE * GAM_PREVALENCE * INCIDENCE_CORRECTION)
print(f"ratio without the coverage target in the denominator: {without_target:.0%}")
```

```
1100 / (21500 * 0.149 * 1.6)
```

Retirez la cible et le même programme affiche la moitié du chiffre. Rien du programme n'a changé.

8.3 Ce qu'exige la couverture

La définition n'est pas contestée. La couverture est ceci.

```
children with SAM currently in the programme
-----
children with SAM in the population
```

et le dénominateur ne peut venir que de la **recherche d'enfants malnutris dans la population qui ne sont pas dans le programme**. Aucun registre ne peut le fournir, puisqu'un registre ne contient que les enfants qui sont venus.

C'est pourquoi les enquêtes de couverture existent comme méthodologie distincte. Les standards actuels — SQUEAC, SLEAC et leurs successeurs — font exactement cela : recherche active de cas dans la communauté, puis classement de chaque cas trouvé comme dans ou hors du programme.

```
# What a coverage survey produces, and a register never can
cases_found = 84
cases_in_programme = 39
print(f"point coverage: {cases_in_programme / cases_found:.0%}")
```

```
c(point_coverage = 39 / 84)
```

Deux mesures de couverture sont en usage et elles répondent à des questions différentes.

- **La couverture ponctuelle** — cas actuellement dans le programme, sur cas trouvés. Convient à un programme aux épisodes de traitement courts.
- **La couverture de période** — cas dans le programme ou récemment déchargés guéris, sur cas trouvés. Convient aux épisodes plus longs et crédite le programme des enfants déjà guéris.

Dites laquelle. Elles diffèrent de plusieurs points et les deux s'appellent couverture.

8.4 La question des barrières est la moitié utile

L'estimation d'une enquête de couverture est un nombre à intervalle large. Son autre produit est souvent plus actionnable — pour chaque cas trouvé et absent du programme, on demande à l'accompagnant pourquoi.

Les réponses se regroupent et renvoient à des remèdes différents.

- **Ignorait que l'enfant était malnutri** — dépistage et sensibilisation communautaire.
- **Ignorait l'existence du programme** — mobilisation.
- **Distance, coût, temps** — décentralisation, qui est aussi la cause d'abandon de la leçon précédente.
- **Refusé précédemment** — un problème de critère d'admission, et la leçon 4 explique comment un enfant peut être refusé sur une mesure tout en étant malnutri selon une autre.
- **Stigmatisation, ou mauvaise expérience antérieure** — qualité des soins.

```
barriers = pd.DataFrame({
    "barrier": ["Did not know child was malnourished", "Distance",
               "Did not know programme existed", "Previously rejected",
               "Carer unavailable"],
    "cases": [17, 12, 8, 5, 3],
})
barriers["share"] = barriers["cases"] / barriers["cases"].sum()
```

```
barriers <- tibble::tribble(
  ~barrier,      ~cases,
  "unaware child malnourished", 17,
  "distance",                12,
  "unaware of programme",     8,
  "previously rejected",      5,
  "carer unavailable",        3
)
```

Rapportez les barrières même quand l'estimation de couverture est trop imprécise pour être employée. Quarante-cinq entretiens ne peuvent pas fixer une couverture à cinq points près et peuvent parfaitement dire que la moitié des cas manqués relèvent d'un problème de dépistage plutôt que d'accès, et ces deux-là n'ont pas le même budget.

8.5 Quoi rapporter sans enquête de couverture

C'est-à-dire la plupart du temps, car une enquête de couverture coûte un argent dont le programme ne dispose généralement pas. Trois options honnêtes, et aucune ne consiste à réétiqueter le rapport.

Rapportez le taux de réalisation du plan, correctement nommé. « Les admissions ont atteint 68 % de la charge planifiée pour la période » est vrai, utile pour l'approvisionnement, et n'affirme rien sur les enfants qui ne sont jamais venus.

Rapportez les admissions face aux cas incidents attendus, hypothèses énoncées. Retirez la cible de couverture du dénominateur, listez les quatre entrées et leurs sources, et étiquetez-le comme une estimation de la portée du programme face au besoin attendu.

Employez le rapport routine-enquête. La dernière leçon du cours DHIS2 a décomposé l'écart entre un registre de dépistage et une enquête ; cet écart est un signal indirect sur qui le programme ne voit pas, et il est gratuit.

```
report = pd.DataFrame({
    "indicator": ["Admissions", "Expected incident cases", "Plan achievement",
                 "Coverage"],
    "value": [1100, 5126, "21%", "not estimated"],
    "note": ["episodes, not children",
            "population x prevalence x incidence correction; three estimates",
            "admissions / expected incident cases",
            "requires active case-finding in the community"],
})
```

```
tibble::tribble(
  ~indicator,          ~value,          ~note,
  "Admissions",        "1,100",        "episodes, not children",
  "Expected incident cases", "5,126",        "three stacked estimates",
  "Plan achievement",   "21%",          "admissions / expected",
  "Coverage",           "not estimated", "requires a coverage survey"
)
```

C'est la dernière ligne qui compte. « **Non estimée** » est un constat, et c'est une bien meilleure ligne de rapport qu'un nombre qui sera cité pendant trois ans comme quelque chose qu'il n'est pas.

8.6 Ce que ce cours vous laisse

Vous savez mesurer un enfant, calculer un score z selon les normes de l'OMS, appliquer les définitions de cas y compris la clause que tout le monde oublie, dire où les deux critères d'admission divergent et pourquoi, produire un rapport de plausibilité SMART et décider si une enquête est croyable, produire une prévalence avec un intervalle et la lire au regard des phases IPC, calculer la performance PCIMA sur le dénominateur que Sphere spécifie, et dire honnêtement ce que votre programme sait et ne sait pas de sa couverture.

C'est le métier de l'analyste nutrition, et c'est le premier des six cours sectoriels du module 4.

Vient ensuite dans le module *Santé publique et épidémiologie appliquée* — incidence, prévalence, cascades et couverture pour le VIH, la tuberculose, le paludisme et la vaccination, ainsi que la courbe épidémique qu'il faudra peut-être tracer dans l'urgence. Il s'exécute sur l'extraction vaccinale que ce programme emploie depuis le module 2, et il démarre là où s'arrêtent les dénominateurs de ce cours — sur les personnes-temps, le dénominateur qui donne son sens à l'incidence.

À propos de cette édition

Ce PDF est produit à partir de la même source que la version web de la formation ; les deux ne peuvent donc pas diverger. L'édition en ligne comporte le code exécutable, les jeux de données téléchargeables et les corrections publiées depuis la production de ce fichier.

Tous les jeux de données cités dans cette formation sont synthétiques ou entièrement anonymisés. Aucun enregistrement ne renvoie à une personne réelle.

Cette formation est publiée en anglais et en français. L'édition française est une adaptation professionnelle reprenant la terminologie employée par l'UNICEF, l'OMS, le PAM, OCHA et les ministères nationaux — et non une traduction littérale.

Consulter et exécuter la formation en ligne sur data-analysis.cassion.dev/fr/cours/nutrition-analysis-cmam-smart

Le contenu pédagogique est diffusé sous licence CC BY 4.0. Vous pouvez l'adapter et le rediffuser, y compris à des fins commerciales, en créditant Cassion.

Généré le 28 juillet 2026.