

L'enquête n'est pas publiable en l'état

Analyse d'enquête SMART et rapport de plausibilité

Organisation	Cluster nutrition commanditaire d'une enquête SMART à 30 grappes
Destinataires	Cluster nutrition et instance nationale de coordination nutritionnelle
Mise à jour	26 juillet 2026
Secteurs	Nutrition
Normes et méthodologies	Enquête SMART · Normes OMS de croissance de l'enfant · Définitions d'indicateurs de l'UNICEF
Données	<code>smart-nutrition-survey-2024</code>

DÉCISION ÉCLAIRÉE

Accepter l'enquête, exclure les données d'une équipe, ou refaire la collecte avant que l'estimation n'oriente une montée en charge.

Tous les jeux de données à l'origine de ce rapport sont synthétiques. Aucun enregistrement ne renvoie à une personne réelle, et aucun chiffre ne décrit un programme réel.

Consulter le projet en ligne sur data-analysis.cassion.dev/fr/projets/smart-survey-analysis

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

Sommaire

Contexte	2
1 Recommandation	2
2 L'estimation, et pourquoi elle ne tranche rien	2
3 Deux défauts de mesure, dans deux équipes différentes	2
4 L'évaluation de plausibilité	3
5 Les options, chiffrées	3
6 Ce que cela n'établit pas	3
7 Prochaines étapes recommandées	4
À propos de cette édition	4

Contexte

Le problème

L'enquête revient avec une estimation de malnutrition aiguë globale qui déclencherait une montée en charge, mais les grappes d'une équipe de mesure sont nettement plus dégradées que les autres. Personne ne pouvait dire s'il s'agit d'une différence géographique réelle ou d'une balance mal réglée, et publier la mauvaise réponse revient soit à gaspiller une réponse, soit à en manquer une.

1 Recommandation

Ne pas accepter l'enquête telle que mesurée. Remesurer les huit grappes de l'équipe 3 si les équipes sont encore sur le terrain. Si elles ne le sont plus, publier en plaçant la comparaison entre équipes dans le corps du rapport plutôt qu'en annexe, et énoncer que l'estimation n'est pas utilisable pour une décision de seuil.

L'enquête a été commandée pour établir si le seuil d'urgence est franchi. Elle ne peut pas répondre à cette question dans son état actuel, et la raison est un défaut de mesure et non un constat nutritionnel.

2 L'estimation, et pourquoi elle ne tranche rien

	Valeur
Enfants analysables	872 sur 930
MAG (Z poids-taille < -2 ou œdèmes)	14,9 %
IC 95 %	11,3 – 18,5 %
MAS	3,9 %
Effet de grappe	2,28

Le seuil d'urgence qui déclencherait une montée en charge est de 15 %. **L'estimation ponctuelle se situe juste en dessous et l'intervalle le franchit.** Sur cette seule base, l'enquête ne peut pas dire si le critère est rempli.

L'intervalle est calculé entre grappes et non entre enfants, car les grappes sont l'unité de sondage. Traiter les 872 enfants comme un échantillon aléatoire simple donnerait un intervalle environ 1,5 fois trop étroit — et placerait tout l'intervalle sous le seuil, produisant une réponse assurée à la mauvaise question.

La méthode complète, y compris la règle longueur/taille de l'OMS et l'extension linéaire au-delà de ± 3 ET, figure dans docs/plausibility-report.md.

3 Deux défauts de mesure, dans deux équipes différentes

Équipe	Enfants	Z moyen	Taille moy.	Poids moy.	MAG
1	208	-0,687	76,11 cm	9,07 kg	15,9 %
2	232	-0,427	75,97 cm	9,26 kg	11,2 %
3	230	-1,099	75,24 cm	8,59 kg	22,2 %
4	202	-0,457	75,41 cm	9,04 kg	9,9 %

C'est l'équipe 3 qui déplace l'estimation. Son score Z moyen se situe deux tiers d'écart-type en dessous des autres, et elle pèse les enfants un demi-kilo plus légers à taille moyenne presque identique — la signature d'une balance mal réglée ou dérivante plutôt que

d'une population plus fragile. Les grappes ont été attribuées aux équipes indépendamment de l'état nutritionnel : un écart réel de cette ampleur serait extraordinaire.

SMART considère un écart entre équipes supérieur à environ 0,3 z comme un problème de supervision. L'écart de cette enquête est de **0,67**.

L'équipe 2 présente un défaut différent, qui ne déplace pas l'estimation. 69,4 % de ses mesures de taille se terminent par ,0 ou ,5, contre 20 % attendus et contre 17,6 à 23,4 % pour toutes les autres équipes. C'est une équipe qui lit la toise au demi-centimètre le plus proche. Cela ajoute du bruit plutôt qu'un biais — le score Z moyen de l'équipe 2 est le deuxième plus élevé des quatre — mais cela gonfle l'écart-type agrégé et constitue de toute façon un constat de formation.

4 L'évaluation de plausibilité

Test	Valeur	Résultat
Mesures impossibles	1,5 % exclues	réussite
Enregistrements signalés (SMART, ± 3 ET)	0,3 %	réussite
Écart-type du Z poids-taille	1,23	alerte
Préférence de chiffre	équipe 2 à 69 % sur ,0/,5	échec
Attraction des âges	24 % sur des années entières	alerte
Sex-ratio	1,01	réussite
Biais entre équipes	0,67 z d'écart	échec

L'alerte sur l'écart-type s'explique par les deux échecs : une équipe mesurant systématiquement bas élargit la distribution agrégée.

L'attraction des âges compte ici pour une raison précise. 23,6 % des enfants sont enregistrés à une année entière exacte, avec des amas à 24, 36 et 48 mois — et **la règle longueur/taille bascule exactement à 24 mois**. Un enfant attiré vers 24 mois est comparé au standard de taille alors qu'il a peut-être 22 mois.

5 Les options, chiffrées

Option	MAG	IC 95 %	n	Franchit 15 %
Accepter tel que mesuré	14,9 %	11,3 – 18,5 %	872	oui
Exclure l'équipe 3	12,3 %	8,6 – 16,0 %	642	oui
Remesurer les grappes de l'équipe 3	—	—	—	—

L'exclusion de l'équipe 3 est rapportée comme analyse de sensibilité, non comme la réponse. Retirer 8 grappes sur 30 modifie la base de sondage, et une prévalence calculée sur le reste n'est plus l'enquête qui avait été conçue. Noter en outre que l'intervalle franchit toujours le seuil : l'exclusion ne résout donc même pas la question qu'elle était censée résoudre.

L'écart entre les deux est de **2,6 points**, et la décision de montée en charge se joue précisément sur cet écart. C'est pourquoi la question de mesure doit être tranchée avant toute utilisation du chiffre.

6 Ce que cela n'établit pas

Quelle équipe a raison. La comparaison montre qu'une équipe diffère ; elle ne montre pas que les trois autres sont justes. Si un exercice de standardisation a été mené avant la

collecte, ses résultats tranchent, et cette analyse ne les contient pas.

Si les grappes de l'équipe 3 sont réellement plus dégradées. L'attribution indépendante rend cela improbable sans le rendre impossible. Seule une nouvelle mesure sépare les deux hypothèses.

La cause de l'un ou l'autre défaut. Une visite de supervision y répond ; un jeu de données non.

7 Prochaines étapes recommandées

1. **Établir si les équipes sont encore sur le terrain.** Si c'est le cas, remesurer les huit grappes de l'équipe 3 est la seule option produisant une estimation publiable.
2. **Récupérer les résultats de standardisation d'avant collecte** s'ils existent. Ils trancheraient la question de l'équipe à croire et ne coûtent rien à obtenir.
3. **Reformer l'équipe 2 à la mesure de la taille**, quel que soit le sort de l'estimation. Une préférence de chiffre à ce niveau se reproduira à la prochaine enquête.
4. **Publier la section « Deux défauts de mesure » dans le corps de tout rapport** portant le chiffre du district. Un lecteur à qui l'on ne montre que 14,9 % ne peut pas savoir qu'un quart de l'échantillon a été mesuré par quelqu'un dont les résultats ne concordent avec ceux de personne.

À propos de cette édition

Ce rapport est produit à partir de la même source que la page du projet ; les deux ne peuvent donc pas diverger. L'édition en ligne comporte le notebook d'analyse, le jeu de données et les corrections publiées depuis la production de ce fichier.

Consulter le projet en ligne sur data-analysis.cassion.dev/fr/projets/smart-survey-analysis

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

Produit le 26 juillet 2026.