

Three answers to one question

Acute malnutrition is 8.7% in the routine screening register, 11.4% in a household survey using the same measure, and 14.9% in a SMART survey using a different one. Decompose the gap into selection and measure, and neither source is wrong.

Pierre Robentz Cassion

Lesson 7 of 8

Routine Data and DHIS2 — Reading it against something else

What this lesson covers

- The meeting
- The three sources
- Hold the measure, vary the selection: 8.7% to 11.4%
- Hold the selection, vary the measure: 11.4% to 14.9%
- What each source is actually good for
- Triangulation, not reconciliation
- The sentence in the report
- When you only have routine data
- What comes next

The meeting

- The nutrition cluster has three figures for acute malnutrition in the same population and the same year.

The three sources — In Python

```
import pandas as pd

muac = pd.read_csv("muac-screening-artibonite-2024.v1.csv")
measured = muac[muac["muac_mm"].notna()]
routine = ((measured["muac_mm"] < 125) | (measured["oedema"] == True)).mean()

print(f"routine screening register: {routine:.1%} on n={len(measured):,}")
```

The three sources — In R

```
library(dplyr)
```

```
muac |>
```

```
  filter(!is.na(muac_mm)) |>
```

```
  summarise(rate = mean(muac_mm < 125 | oedema), n = n())
```

The three sources

Source	Measure	Population	Estimate	n
Routine screening register	MUAC	Children brought to screening	8.7%	4,146
Household survey	MUAC	Population sample, weighted	11.4%	648
SMART survey	Weight-for-height z	Population sample	14.9%	874

Hold the measure, vary the selection: 8.7% to 11.4%

- Routine data measures the served population, not the population — That is not a defect of the system; it is what a . . .

Hold the selection, vary the measure: 11.4% to 14.9% — In Python

```
comparison = pd.DataFrame({
    "source": ["Routine register", "Household survey", "SMART survey"],
    "measure": ["MUAC", "MUAC", "weight-for-height z"],
    "population": ["screened", "population sample", "population sample"],
    "estimate": [0.087, 0.114, 0.149],
})
comparison["vs_previous"] = comparison["estimate"].diff()
print(comparison)
```

Hold the selection, vary the measure: 11.4% to 14.9% — In R

```
tibble::tribble(  
  ~source,      ~measure, ~population,      ~estimate,  
  "Routine register", "MUAC",  "screened",      0.087,  
  "Household survey", "MUAC",  "population sample", 0.114,  
  "SMART survey",    "WHZ",   "population sample", 0.149  
) |> mutate(vs_previous = estimate - lag(estimate))
```

Hold the selection, vary the measure: 11.4% to 14.9%

- Selection accounts for 2.7 points. Measure accounts for 3.5 — Between them they account for the whole of the 6.2-point. . .

What each source is actually good for

Source	Good for	Not for
Routine register	Caseload, workload, supply planning, trend within the served population	Prevalence, coverage denominators
Household survey	Population prevalence with an interval, equity across strata	Anything monthly; anything about individual facilities
SMART survey	Prevalence against the IPC and emergency thresholds, comparability with other surveys	Anything more often than annually

Triangulation, not reconciliation

- **A widening gap between routine and survey** means screening coverage is falling, or the hardest-to-reach are becoming. . .
- **A routine figure that moves while the survey does not** is usually about case finding — a new outreach round, a new. . .
- **A survey figure that moves while routine does not** means the programme is not seeing the change. That is the most. . .

Triangulation, not reconciliation — In Python

```
def gap_over_time(routine_series, survey_series):  
    return pd.DataFrame({  
        "routine": routine_series,  
        "survey": survey_series,  
        "ratio": survey_series / routine_series,  
    })
```

Triangulation, not reconciliation — In R

```
gap_over_time <- function(routine, survey) {  
  tibble::tibble(routine, survey, ratio = survey / routine)  
}
```

Triangulation, not reconciliation

- Track the ratio, not the difference — A ratio of survey to routine is roughly the inverse of screening coverage, and it. . .

The sentence in the report — Example

Acute malnutrition in the district, 2024:

14.9% (95% CI 12.3-17.9) by weight-for-height z-score, from a 30-cluster SMART survey of 874 children, design effect 1.5.

11.4% by MUAC in the same population, from a household survey of 648 measured children, weighted to the sampling frame.

8.7% by MUAC among 4,146 children brought to routine screening. This is a measure of the screened population, not a prevalence estimate, and is reported here for comparison with previous rounds of screening only.

The gap between the first two figures is a measurement difference; between the second and third, a difference in who was measured. Neither indicates an error.

When you only have routine data

- **Report caseload, not prevalence.** "3,700 children screened, 362 acutely malnourished by MUAC" is a true sentence. . .
- **Report the trend, not the level.** A routine series is far more reliable about direction than about level, provided. . .
- **Anchor to the last survey.** Where a survey exists, the routine-to-survey ratio from that year can be carried forward. . .

What comes next

- Every lesson in this course has been a version of one question.

Read the full lesson, with runnable code

[https://data-analysis.cassion.dev/courses/routine-data-and-dhis2/
dhis2-07-routine-versus-survey](https://data-analysis.cassion.dev/courses/routine-data-and-dhis2/dhis2-07-routine-versus-survey)