

4.2

Measuring the problem: Basic statistics

Authors

Christopher Garimoi Orach, Makerere University, School of Public Health, Kampala, Uganda.

Ngoy Nsenga, WHO Regional Office for Africa, Nairobi, Kenya.

Olushayo Olu, WHO Country Office, Juba, South Sudan.

Megan Harris, Public Health Wales, Swansea, Wales.

4.2.1 Learning objectives

To understand the following in the context of Health EDM:

1. Basic statistical concepts.
2. Epidemiologic study designs.
3. Commonly used sampling methods.
4. Estimation of sample size.

4.2.2 Introduction

Statistics are used to describe the health status of population groups, quantify disease burden and estimate the effects of interventions. This is especially important in Health EDM, where health authorities making decisions about the use of limited resources need to be able to identify the best possible programmes for prevention and care so that they can prioritize key interventions. One of the prerequisites of data analysis is to collect data that will allow the research questions to be answered and hypotheses to be tested (Chapter 3.5). The kind of statistical analyses chosen will depend on the type of data that were collected through research, routine data collection or surveillance data.

Case study 4.2.1 provides an example of how the data collection for statistics was conducted in humanitarian settings.

4.2

Case study 4.2.1

Measuring the public health problem in a human-made disaster in Sub-Saharan Africa

An armed conflict in Sub-Saharan Africa resulted in a major humanitarian crisis. The conflict internally displaced more than one million people into camps which were largely managed by the United Nations (1). Camps for internally displaced persons may have poor living conditions, overcrowding and inadequate access to social services that predispose the displaced populations to outbreaks of infectious diseases such as measles, cholera, malaria, and hepatitis E (2). The Early Warning Alert and Response System (EWARS) was established to address the need for good quality and real-time data for timely detection and response to epidemics in support of the Early Warning Alert and Response Network (EWARN) (3), a system that supports surveillance and response in humanitarian settings where routine systems are unavailable or underperforming (4).

The system collects real-time data on infectious diseases, injuries, trauma and nutrition from health facilities managed by frontline health partners in the camps and conflict-affected areas. Data are entered at the facility level and automatically uploaded into a central database. Automated analysis is conducted, a weekly bulletin is generated and disseminated to all health partners on a regular basis. The system resulted in drastic improvements in the timeliness (69%) and completeness (73%) of reporting from the camps and conflict-affected locations and timely detection of several outbreaks including the cholera epidemic of 2016 and measles outbreaks of 2018 to 2019 (5). The system also provides detailed case-based and laboratory data which are used for better characterization and response to outbreaks and for research purposes. Furthermore, the system contributes to improvements in the national Integrated Diseases Surveillance and Response System and has been expanded to generate monthly information on health service functionality and nutrition status. Poor mobile network coverage in the conflict-affected areas of the country remains a key challenge as data are transmitted electronically.

The EWARS has proven to be a good tool in the generation of data for public health decision making during humanitarian crises while also serving as foundation for strengthening disease surveillance during the transition from humanitarian to development programming. The system is also a major repository of secondary research data.

4.2.3 Types of quantitative data

The two main types of quantitative data are categorical and continuous. Categorical data can be either dichotomous (taking only one of two possible values) or polytomous (having more than two distinct categories). Dichotomous data are considered binary – for example, vital status might be either alive or dead, a community might have either been exposed or not exposed to a toxic spill and someone might have either received or not received an intervention. Polytomous data have more than two categories and have a number of different attributes. It may be ordinal, being rank-ordered, typically based on a numerical scale that is comprised of a small set of discrete classes or integers, but may not always have a specific set interval between integers (for example, socio-economic status or income level). Alternatively, the categories might not be in any order (for example, types of injury or cause of death).

Continuous data are measured on a continuum and, theoretically at least, can have any numeric value over a continuous range, with the level of granularity dependent on the precision of the measurement instrument. Interval data are a form of continuous data in which equal intervals represent equal differences in the property being measured, for example temperature. Ratio data are another form of continuous data, which have the same properties as interval data, plus a true definition of an absolute zero point – for example weight or height (6).

4.2.4 Types of statistical analysis

Statistical methods can be divided into two main branches: descriptive and inferential. Descriptive statistics are commonly used to categorize, display and summarize data; inferential statistics are used to make predictions based on a sample obtained from a population or some large body of information. These inferences can be used to test specific research hypotheses (7). This chapter covers the basic statistical principles that should be considered when choosing a study design and conducting the study. It includes examples and definitions of issues such as summary statistics and the calculation of the sample size needed for a study. Other chapters in this book deal with the development of the research question for a study (Chapter 3.5), study design (Chapter 4.1) and data collection (Chapter 4.4); more advanced statistical techniques are covered in Chapter 4.5.

4.2.5 Descriptive statistics

Descriptive statistics are typically used simply to calculate, describe and summarize the collected data in a logical, meaningful, and efficient way. Descriptive statistics do not allow any conclusions to be drawn regarding the validity of research hypotheses. They might include measures of central tendency (such as the mean, the median and the mode) to show the most representative value of the data set. They are usually accompanied by a measure of dispersion (such as the standard deviation or inter-quartile range) to indicate the degree of variation of values within a data set or the level of dispersion of observations around the measure of central tendency. Some of these are described below.



4.2

Measures of central tendency

Mean: the mean (sometimes referred to as the arithmetic mean) is the most common measure of central tendency. It is calculated by dividing the sum total of all observations by the number of records. One advantage of the mean is that, because its calculation includes the summing of all the observations, its value takes into account all the data. However, this characteristic of the mean also makes it especially sensitive to extreme values among the observations, which can skew this central tendency towards extreme outliers. Thus, the mean can be a misleading measure if the data set contains such outliers.

Median: this is the observation that divides the distribution into two equal parts. In other words, when all observations are ranked from the lowest to the highest, the median is the observation that is located at the half way point. Therefore, the median can only be determined for observations that are ranked by value or size and is less influenced by extreme values. The median can be used to compare groups on certain characteristics (for example, to compare the age between two groups of children or to compare number of days of exposure to extreme weather for people in different regions).

Mode: this is the observation or value that appears most frequently in a set of data. The mode is identified by noting the observation that occurs the most or value that has the highest number of records. The mode has the advantage of being easy to identify by simply counting the frequency of the records presenting that value. However, its main disadvantage is its potential lack of stability as a measure of central tendency because it can change if the data set is categorized or even defined in different ways. The mode can be used to determine, for instance, which socioeconomic group has the highest number of individuals.

Measures of dispersion

Standard deviation: this is the square root of the deviance, which is calculated by squaring and summing the difference between each observation and the arithmetic mean. The sum is then divided by the total number of observations. In the same population, the standard deviation is more stable from one sample to another. When comparing two groups or samples, a group or sample with a relatively smaller standard deviation indicates that the members of this group are more homogenous (or similar to each other) than the group with a large standard deviation. If the observations in a data set have a normal distribution, 70% of observations will lie within one standard deviation of the mean and 95% within two standard deviations (8).

Standard error: This measures the amount of variance in the sample mean and is calculated by dividing the standard deviation by the square root of the number of observations in the sample. The standard error is used to indicate how well the true population mean is likely to be estimated by the sample mean.

Range: This represents the difference between the highest and the lowest values of the distribution and can be used to give complementary information to other statistics, such as the mean. When two distributions seem to have similar means, the range can provide an additional layer of information to distinguish the characteristics of the two distributions.

However, one important disadvantage of the range is that it will be influenced by extreme values. This means that a change in a single record that was the highest or lowest value could have a substantial impact on the range. The range can also be expressed in quartiles or in percentiles to show the highest and lowest values in different parts of the distribution (such as the range of ages for children and for adults in a sample).

Interquartile range: Just as for calculating the median as the half-way point in a series of observations, the interquartile range requires the observations to be ranked from the lowest to the highest. The interquartile range median is then the difference between the lower (25th percentile) and the higher (75th percentile) quarters of the observations.

Confidence interval: This is derived from the standard error of the mean. The confidence interval (usually 95%) shows the range within which the true population value is likely to fall, based on the sample statistical values and probability data distributions.

4.2.6 Inferential statistics

In the context of research into the effects of interventions (as discussed in Chapters 4.1 and 4.3), inferential statistics allow researchers to make a valid estimate of the association between an intervention and its effect in a specific population, based upon their representative sample data. Inferential statistics allow researchers to make generalizations or inferences from the results obtained from the sample to the populations from which the samples were drawn. Approaches to inferential statistics include the estimation of parameters, and the testing of research hypotheses. Inferential statistics vary depending on the type of statistical tests applied in the analysis. For instance, they might use correlation coefficients to assess the correlation and association between risk factors and outcome, or use an odds ratio to measure the probability of an event occurring.

4.2.7 Rapid needs assessments

Rapid needs assessments (as also discussed in Chapter 2.1) will usually require basic statistical analyses to be conducted. For instance, in disaster settings, rapid needs assessments often use survey sampling techniques in the field to rapidly determine the health status and basic needs of an affected community. Emergency response requires immediate information on health status and community needs. Such information must be gathered and analysed quickly. In many cases, an assessment may need to be initiated and completed within 72 hours. Speed is critical because circumstances can change dramatically with time, and outdated information may therefore be of little use to response personnel (9). However, these surveys need to be conducted in a statistically robust and valid manner to support decisions about the response. Various areas of consideration (such as disease states or conditions) might need to be measured using various statistical parameters – such as prevalence, incidence and attack rate (see below).

A rapid health needs assessment is often carried out at a single point in time, using a cross-sectional study design. Key stakeholders should be involved in the survey process, and it is important to identify specific



4.2

targeted groups as the study population, depending on the objective of such needs assessment. For example, when undertaking a nutrition assessment, the study population may include all children under the age of 5 years and their parents. The sample size for the study (see below) might not be estimated statistically but may simply be based on the population who are being studied.

Rapid needs assessments will collect data on the population and may include the number of displaced or affected people and their demographic characteristics (for example, the number of women, men, children, pregnant women and persons with disability). It might also be important to collect data on the proportion of people with shelter, in order to establish the shortfall in shelter requirements for the displaced population (such as refugees or internally displaced persons). Data should also be collected on the available resources (see also Chapter 3.1), including health systems. This might include the number and type of health facilities, number and category of health workers and types of health services available. Depending on the situation, data may also be collected from other sectors such as water and sanitation, education, food security, protection and so on. It might also be gathered to establish a picture of other baseline features, such as numbers of medical staff still working per 1000 people in the population, vaccination rate for key vaccines or rate of severe acute malnutrition. During emergencies, the values of these indicators are usually compared to reference values and norms, such as the Sphere standard to evaluate the status of population humanitarian condition (10). There is more information on health indicators in Chapter 2.2.

4.2.8 Epidemiologic Measures

This section provides a brief review of some key terms used in epidemiology to describe data about diseases.

Population

In the epidemiology of disasters (Chapter 2.1), the definition of the “population” can vary depending on the situation. In general, the term is used to refer to people living in a defined area, such as a refugee camp, settlement, village or neighbourhood. However, in some situations, it may refer to groups of people being affected by an emergency, who do not necessarily live in a well-defined area. For instance, in an infectious disease outbreak, population may refer to groups of people with a specific characteristic, such as a profession, lifestyle or activity that predisposes them to the disease (for example, farmers, butchers, or those in school settings). It might also be necessary to count subgroups of the population, such as the number of women or the number of children under 5 years of age.

In some cases, the total population figure will be the denominator for calculating health indicators (Chapter 2.2). For example, it might be used to estimate the proportion of people out of the total population who were made homeless after an earthquake, the proportion of pregnant women who are likely to give birth in the days after a disaster, or the proportion of children in an internally displaced person (IDP) camp who have not been vaccinated against measles.

Usually, the census or a registration system might be relied on as the most accurate method of estimating the population. However, in an emergency, it might be necessary to use other methods (Chapter 2.4), such as mapping the IDP camp and dividing it into smaller sections, with the population size of each section estimated using sample surveys.

Depending on the type of data being collected and the context, gathering information from individuals can sometimes be perceived as intrusive. It is, therefore, important to identify and implement methods to count people and cases that maintain the dignity of the individuals involved, using appropriate ethical oversight (Chapter 3.4). This is especially important if public health priorities (speed, accurate information) and human rights priorities (privacy, consent) might come into conflict during data collection.

Data Analysis

Basic data analysis can be used to provide information to guide the development and implementation of operational plans for Health EDRM. The information is often summarized into a minimum set related to person, place and time. Minimum data analysis can generate basic answers to questions such as: who is affected or most at risk? Where are those affected or at most risk? What is the trend of the impact of the events on the target population? Subsidiary, basic analysis can provide insight into major risk factors making the target population vulnerable or rendering them resilient to the effect of the hazard. In addition to the descriptive statistics outlined above, epidemiology uses measures of morbidity and mortality and these rely on the quantification of various aspects of health, outlined below.

Prevalence

This is useful for understanding the overall burden of a disease on a population, since it describes how common a particular condition is at a given point in time (point prevalence) or the existing and new cases that happen over a set period of time, such as 12 months (period prevalence). Prevalence is a calculation of the existing cases and is determined by the rate of new cases occurring, the rate of recovery and the rate of deaths. Prevalence is often used for conditions that are longer lasting or for which an on-set date may be more difficult to recall (for example, the number of people suffering anxiety related to a disaster).

Incidence

This is the number of new cases of the condition occurring in a given population during a defined period of time. There are different ways to calculate incidence, based on the condition, issue or disease. The most common is the cumulative incidence, which is the number of new cases in a specific time period divided by the number of people who were initially disease (or condition) free at the start. For example, if there were 120 new measles cases in one week among 18 000 people in an IDP camp, this would give an incidence rate of 6.7 per 1000 per week. The incidence rate is useful when discussing or comparing acute, communicable diseases of short duration.

Attack rate

This is the cumulative incidence rate of a disease in a specified population over a given period of time. It is usually used during epidemics and is calculated as the percentage of the population with a condition out of the

4.2

whole population (for instance, those with the condition and healthy, susceptible people) (Table 4.2.1). The attack rate can help when calculating the resources needed to respond to an outbreak. It also provides an idea as to the magnitude of an outbreak in a community or a geographic entity. If immunity to the disease (as a result of vaccination or prior infection for instance) is measured, this may allow some of the population to be removed from the denominator.

Table 4.2.1 Example of incidence and attack rate for measles among 18 000 refugees

Week	New cases per week	Weekly Incidence Rate	Attack Rate
1	120	6.7 per 1000	0.67%
2	150	8.3 per 1000	1.50%
3	80	4.4 per 1000	1.94%

Case fatality rate

This is the number of deaths from a specific disease during the observational period, divided by the number of cases of that disease during that period, multiplied by 100 (to calculate a percentage). The case fatality rate is used mainly in infectious diseases, such as cholera, dysentery, malaria and measles. It provides a useful guide to assess the virulence of the disease, its severity and the effectiveness and quality of care.

Mid-interval population

This can be estimated by adding together the number of people in the population at the start of the period of observation and the number at the end, and dividing this by two. Alternatively, it can be calculated as the average size of the population during the period. Population data are usually collected from official government census reports or other administrative documents, such as the birth and deaths registry. It may already be available from national statistical offices and published online.

Benchmarks

These are standards or reference values for indicators that serve as signposts to let the researcher, or other interested people such as policy makers, know what has been achieved or how severe a situation is. They can include key mortality indicators such as the infant mortality rate, cause-specific mortality rate and case fatality rate discussed below.

4.2.9 Demographic indices

Demographic indices include statistics such as fertility rates, birth rates, growth rates and mortality rates.

Crude birth rate

This is calculated as a proportion by dividing the number of live births by the number of people in the mid-interval population, and multiplying the value by 1000 (or other amount depending on the population size) to create a rate.

Crude growth rate

This is the crude birth rate minus the crude mortality rate. It provides information on the growth or decline of a population, in the absence of migration.

Crude mortality rate

This is calculated as a proportion by dividing the number of deaths at all ages by the number of people in the mid-interval population, and multiplying the value by 1000 (for annual or monthly rates) or 10 000 for daily crude mortality rate. This crude rate does not adjust for the age distribution of the population, and should not be used to compare across different populations.

Infant mortality rate

This is calculated by dividing the number of deaths in children under one year of age by the number of live births during the same period and multiplying this by 1000 (or other amounts depending on the population size). Although this is conventionally referred to as a rate, it is really a ratio. This is because in a rate, those counted in the numerator must also be part of the denominator (for example, the number of deaths due to measles divided by cases of measles). However, in the infant mortality rate, some of those children who die during the specified interval (the numerator) might not have been born during the same interval (the denominator).

Cause-specific mortality rate

This is the number of deaths from a specific cause during the observational period divided by the number of people in the mid-interval population (or other denominator of the population), multiplied by 100 to provide a percentage.

Age-specific mortality rate

Because different populations have different characteristics and age structures it is not meaningful to compare the crude mortality rate for different settings or countries. For example, a high proportion of elderly people in a population will give it a high crude mortality rate and, as a result, the crude mortality rate of the Plurinational State of Bolivia and that of the USA may be very different because of the underlying age-distribution rather than the likelihood of an individual dying. To overcome this, age-specific mortality rates are calculated. There are two different methods of standardizing population statistics – direct standardization and indirect standardization. More information on these methods can be found in Gerstmann (11).

4.2.10 Epidemiological Studies

Epidemiological studies can be descriptive, analytical or both. Descriptive studies are used to describe exposure and disease in a population (see Chapter 3.2), and can be used to generate hypotheses, but they are not designed to test hypotheses. Analytical studies are designed to test hypotheses, and are designed to evaluate the association between an exposure or intervention and a disease or other health outcome (see Chapters 4.1 and 4.3).

Epidemiological studies can be cross-sectional, prospective or retrospective. A cross-sectional study is taken at a specific point in time. A

4.2

prospective study is one where the study starts before the exposure and outcomes are measured moving forward in time. A retrospective study is one where the study starts after the exposure has begun and, in some cases, the outcomes have occurred and been measured. It works backwards in time. Epidemiological studies can also be experimental or observational and some of the terminology important for epidemiological studies is described below.

Exposure

This is the risk factor (agent, experience or procedure for example) that is suspected to have caused the disease or condition. In statistical terms, exposure is often called the independent variable.

Outcome

This is the disease, condition or other endpoint being measured. In statistical terms, the outcome is often called the dependent variable.

4.2.11 Descriptive studies

Descriptive studies describe an event, condition or disease state in terms of time, place and person. They include:

- Case series or record review.
- Descriptive incidence study (active surveillance)
(for example, collecting information on all cholera cases, by age, sex, location of hut, nearest water source and duration of stay in an IDP camp).
- Descriptive prevalence study (cross-sectional survey)
(for example, a study of prevalence of acute malnutrition among children under 5 years of age).
- Ecological study (for example, times series analysis of the impact of air pollution on respiratory morbidity and mortality).

4.2.12 Analytical studies

Analytical studies examine the relationship between a possible cause (or exposure or intervention) and its effect (disease or condition). These are generally developed to test a hypothesis, which could have been developed from descriptive studies previously undertaken. Two common examples of analytical studies are cohort studies and case-control studies:

Cohort study

In a cohort study, a population is followed over time (either prospectively or retrospectively). There are usually two study groups: those exposed to a certain exposure – which may be either a risk factor (such as diet deficient in vitamin C) or a protective factor (such as measles immunization) – and those not exposed. The cohort study follows both groups over a period of time and estimates incidence of the outcome in each group. The measure of association in this study design is the relative risk, which is the ratio of the incidence of disease in the exposed group to the incidence of disease in the non-exposed group. Cohort studies can be carried out in many time frames, from days to decades.

Case-control study

In a case-control study, the two groups being compared are people who meet the criteria (or case definition) of the disease or other outcome and people from the same or similar population who do not, as a control group. This retrospective design is used to determine who was exposed to certain factors (contaminated water, for example) and who was not exposed and whether exposure in those who have the outcome is different to those without. The measure of association in this study design is often the odds ratio, which is the ratio of the odds of disease in the exposed group to the odds of disease in the non-exposed group. The odds of disease is the proportion of people with the disease divided by those without it.

4.2.13 Sampling Methods

When choosing the people to include in a study, a variety of sampling methods are available:

Non-probability or judgemental sampling

For example – convenience, snowballing or quota sampling.

Probability sampling

Probability sampling includes simple random sampling, systematic sampling and cluster sampling; Table 4.2.2 shows the advantages and disadvantages of each of these specific methods.

Simple random sampling:

This would lead to a fully random sample by using a method such as a random number table to draw the sample from a whole population to which all the members belong.

Systematic sampling

This involves choosing the first member of the sample of the whole population using a random number and choosing the rest of the sample by proceeding at a fixed interval.

Cluster sampling

This involves the random selection of a cluster (such as a village, school or hospital) and then random sampling of the individuals from within the selected clusters.

Table 4.2.2 Advantages and disadvantages of different types of probability sampling

Type of probability sampling	Advantages	Disadvantages
Simple random sampling	Minimal bias. Every member has an equal chance of being included (which can balance confounding factors).	Must enumerate all members of the population, which is expensive and sometimes not feasible. Can miss geographical clusters (such as people from a minority ethnic group living in one part of an IDP camp).
Systematic sampling	Guarantees a broad geographical representation. Do not have to have prior knowledge of the total number of people who could be selected for the study.	May be expensive and time consuming to ensure full randomization.
Cluster sampling	Easier to conduct, less travel time and cost. Do not need a complete list of the sampling units.	Bias toward more dense areas, such as town centres.

If a sample is used for a study, rather than the whole population, this leads to an estimate of what the results might be for the population as a whole. If a series of samples is taken, these are likely to give different values, but providing the samples have been selected correctly there should be little variation between them. However, in order to provide an estimate of this variation, confidence intervals are often used to show the extent of the variation. The confidence intervals provide the upper and lower limits of this range. For example, if the mean for a sample was 12% and the standard deviation was 2%, the 95% confidence interval would be shown as 10 to 14%.

4.2.14 Sample size calculation

If it were possible for a research study to include the whole population of interest, sampling would not be necessary, but covering a whole population would usually require too much money, time or personnel. Therefore, researchers need to rely on a population subset: the sample. This allows them to seek reasonably valid answers to their research questions, but they first need to estimate the size of the sample needed to achieve this. Determining the appropriate sample size for a study is a fundamental aspect of all research; this is because having an adequately-sized sample ensures that the information the study yields will be reliable, regardless of whether the data ultimately suggest an important difference between the impact of a disaster on different types of people, or the effects of intervention and control in a randomized trial.

Two types of false conclusion may occur when inferences about the whole population are derived from a study of a sample of the population. These are called Type 1 and Type 2 errors, whose probabilities are denoted by the symbols α and β . A Type 1 error occurs when one concludes that a difference exists between the groups being compared when, in reality, it

does not. This is akin to a false positive result. A Type 2 error occurs when one concludes that a difference does not exist when, in reality, a difference does exist, and it is equal to or larger than the effect size defined by the alternative to the null hypothesis (12).

The calculation of a sample size for a research study depends on the type of study being planned, the data to be collected, the outcomes being measured and the hypothesis being tested (13). More information is available in the texts listed in the further reading section (4.2.17) but, in general, sample size estimation depends on the level of confidence and precision. The following formula can be used to calculate the sample size for a binary outcome:

$$n = \frac{Z^2 pq}{d^2}$$

n corresponds to the sample size in each of the groups; Z is the level of confidence chosen (95% confidence, $Z = 1.96$; 90% confidence: $Z = 1.68$); g is the design effect and a usual value for this situation is 2; p is expected proportion of the population with the characteristic of interest (such as acute malnutrition), q is $1-p$; and d is the precision (in proportion of one; if 5%, $d = 0.05$).

This formula shows that in order to increase the level of confidence or precision, the sample size must be increased. Therefore, when a study is trying to detect a small effect with high precision (such that the entire width of confidence interval would be consistent with a beneficial effect of an intervention, for example), the study will need to be much larger than when the study is testing a hypotheses that there is a large effect.

4.2.15 Conclusions

This chapter presents an introduction to basic statistical concepts, epidemiologic study designs, commonly used sampling methods and estimation of sample size. It provides basic statistical knowledge to support effective Health EDRM.

4.2

4.2.16 Key messages

- o **Statistical analyses of quantitative data from research studies and the results these generate are vital to a variety of types of research in Health EDRM. They help by estimating disease burden (to help with the distribution of humanitarian assistance, for instance), the health consequences of disasters for populations (to help with planning for future needs, for example) and the effects of interventions, actions and strategies (to prioritize the elements to include in humanitarian assistance, for example). They often require the contribution of partners with diverse disciplines.**
- o **Practitioners need to understand a variety of methods of data collection and analysis, and apply those most relevant to their research question if they are to answer it reliably. This might include surveys, cohort studies, case control studies or experimental studies such as randomized trials for quantitative research and the use of qualitative methods where appropriate.**
- o **Research in emergency settings is constrained by ethical concerns (Chapter 3.4) and limited resources, increasing both the challenges of conducting rigorous epidemiological research and the importance of reliable statistical analysis of the data that are available.**

4.2.17 Further reading

Gerstman B. Basic Biostatistics: Statistics for Public Health Practice (2nd edition). Burlington, MA: Jones & Bartlett Learning. 2014.

Horney JA . Disaster Epidemiology: Methods and Applications. London, UK: Elsevier. 2017.

Ricci EM, Pretto EA. Disaster Evaluation Research: A field guide. Oxford, UK: Oxford University Press. 2019.

4.2.18 References

1. Global Emergency Overview Snapshot, 6 - 12 May 2015 – World. ReliefWeb. 2015. <https://reliefweb.int/report/world/global-emergency-overview-snapshot-6-12-may-2015> (accessed 12 June 2019).
2. Outbreak surveillance and response in humanitarian emergencies. WHO. 2012. http://www.who.int/diseasecontrol_emergencies/publications/who_hse_epr_dce_2012.1/en/ (accessed 13 April 2018).
3. South Sudan health crisis worsens. WHO. 2016. <https://www.who.int/news-room/feature-stories/detail/south-sudan-health-crisis-worsens-as-more-partners-pull-out-and-number-of-displaced-rises> (accessed 12 June 2019).
4. South Sudan weekly disease surveillance bulletin 2019. WHO Regional Office for Africa. 2019. <https://www.afro.who.int/publications/south-sudan-weekly-disease-surveillance-bulletin-2019> (accessed 12 June 2019).
5. WHO and MoH. South Sudan (EWARN) Early warning and disease surveillance bulletin. 2015. <http://www.who.int/hac/crises/ssd/epi/en/index3.html> (accessed 22 June 2017).
6. Vetter TR. Fundamentals of Research Data and Variables: The Devil Is in the Details. *Anesthesia and Analgesia*. 2017; 125: 1375-80.
7. Overholser BR, Sowinski KM. Biostatistics primer: part I. Nutrition in Clinical Practice. 2007; 22: 629-35.
8. Kirkwood BR, Sterne JAC. *Essential Medical Statistics* (2nd edition). Malden MS: Blackwell Science. 2003.
9. Malilay J, Heumann M, Perrotta D, Wolkin AF, Schnall AH, Podgornik MN, et al. The role of applied epidemiology methods in the disaster management cycle. *American Journal of Public Health* 2014; 104: 2092-102.
10. Sphere. *The Sphere Handbook: Humanitarian Charter and Minimum Standards in Humanitarian Response* (4th edition). Geneva, Switzerland. 2018. <https://handbook.spherestandards.org/en/sphere/#ch001> (accessed 16 January 2020).
11. Gerstmann B. *Basic Biostatistics: Statistics for Public Health Practice* (2nd edition). Burlington, MA: Jones & Bartlett Learning. 2014.
12. Hazra A, Gogtay N. Biostatistics Series Module 5: Determining Sample Size. *Indian Journal of Dermatology*. 2016; 61: 496-504.
13. Devane D, Begley CM, Clarke M. How many do I need? Basic principles of sample size estimation. *Journal of Advanced Nursing*. 2004; 47(3): 297-302.