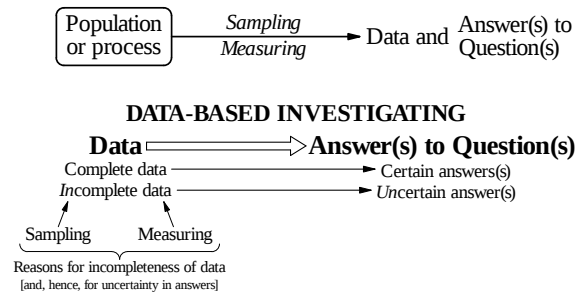


RESPONSE MODELS IN STAT 231: Modelling Sample Error and Measurement Error

As summarized in the two schemas at the right, statistics is concerned with *data-based investigating* of some population or process to *answer* one or more (statistical) *questions* of interest.

- * If the investigating yields *complete* information, we can obtain a *certain* answer; that is, an answer we can *know* is correct.
- * If the investigating yields *incomplete* information, we *cannot* know an answer is correct (an *uncertain* answer) – in fact, it is *unlikely* a *numerical* answer (an ‘estimate’) is (exactly) correct;
 - sampling and measuring yield data (and, hence, information) that are *inherently* incomplete.

In this Highlight #74, we develop a probability a model for the investigative processes of sampling (which involves (selecting and estimating) and measuring; this model allows us (e.g., in Statistical Highlight #2) to quantify the *likely* size of sample error and measurement error – that is, to quantify uncertainty from these two sources – for numerical answers to some types of questions.



1. Ideas from Probability

To go from probability to statistics, ideas we need from probability modelling (e.g., in STAT 230) are as follows.

- * Using the normal distribution to model the shape of appropriate data distributions, as summarized by the probability statement (HL74.1) about the random variable Y at the right. $Y \sim N(\mu, \sigma)$ -----(HL74.1)
- * A random variable which is a sum or an average of (probabilistically independent) normal random variables also has a normal distribution and its mean and standard deviation can be expressed in terms of the mean(s) and standard deviation(s) of the random variables that make up the sum or average. For example, for n probabilistically independent $N(\mu, \sigma)$ random variables $Y_j, j = 1, 2, \dots, n$, as given in equations (HL74.2) and (HL74.3) at the right below:
 - their sum (T) [or *total*] has a normal distribution with a *larger* mean and standard deviation than the individual random variables by respective factors of n and $\sqrt{n}$. $T \sim N(n\mu, \sigma\sqrt{n})$ -----(HL74.2)
 - their *average* ($\bar{Y}$) has a normal distribution with the *same* mean μ as the individual random variables but a standard deviation that is *smaller* by a factor of $1/\sqrt{n}$. $\bar{Y} \sim N(\mu, \sigma/\sqrt{n})$ -----(HL74.3)

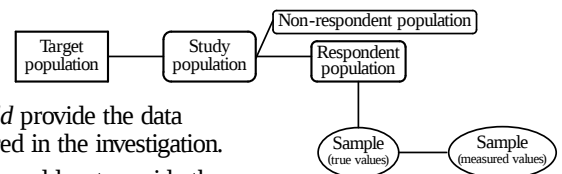
If the requirement for *normality* of the distribution of the individual random variables in a total or an average is relaxed, the Central Limit Theorem (CLT) approximation allows us to write equations (HL74.4) and (HL74.5); the *accuracy* of this approximate normality of a total or an average improves as:

- the value of n gets *larger*;
 - the distribution(s) of the Y_j s get *more symmetrical*.
- $$T \div N(n\mu, \sigma\sqrt{n}) \text{ -----(HL74.4)}$$
- $$\bar{Y} \div N(\mu, \sigma/\sqrt{n}) \text{ -----(HL74.5)}$$

2. Ideas from Statistics

The schema at the right reminds us that data-based investigating in statistics is concerned with five groups of elements or units and their **attributes**: a quantity defined as a function of response (and, perhaps, explanatory) variates over the group (e.g., an average).

- * **Target population**: the group of elements to which the investigator(s) want Answer(s) to the Question(s) to apply.
- * **Study population**: a group of elements *available* to an investigation.
- * **Respondent population**: those elements of the study population that *would* provide the data requested under the incentives for response offered in the investigation.
- * **Non-respondent population**: those elements of the study population that *would not* provide the data requested under the incentives for response offered in the investigation.
- * **Sample**: the group of units selected from the respondent population *actually used* in an investigation – the sample is a *subset* of the respondent population (as indicated by the vertical line in the schema above).

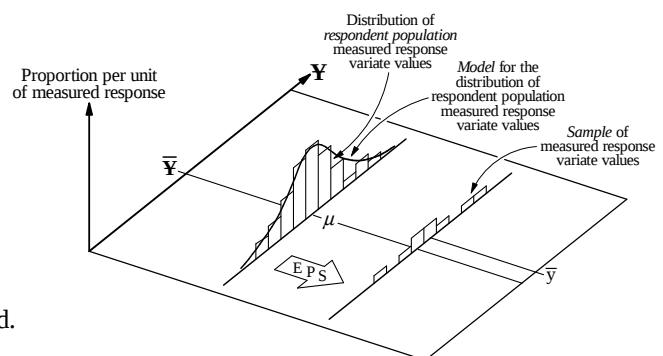


As indicated pictorially in the 3-dimensional diagram at the right, we think of a *respondent* population as having a *distribution* of values for some measured *response* variate (Y) of interest; this distribution can be displayed as a *histogram*. Simple population *attributes* of interest are:

- * Size: N [the number of *elements*];
- * Location: $\bar{Y}$ [the *average*];
- * Variation: S [the (*data*) *standard deviation*];

Unless N is to be estimated, it is usually considered to be:

- known (or ignored), and
- large compared to the size (n) of any *sample* (to be) selected.



(continued overleaf)

The probability model (HL74.1) for the shape of the histogram is included in the diagram overleaf on page HL74.1 at the bottom right, which also reminds us that μ , the mean of the model, represents the respondent population average, $\bar{Y}$, of the response variate. The *sample* average is $\bar{y}$ and equation (HL74.3) is a model for the corresponding random variable $\bar{Y}$; the method of *selecting* the sample [equiprobable selecting (EPS)] is (part of) the reason we can treat the measured response variate values, y_j , $j = 1, 2, \dots, n$, of the n units in the sample as values, y_j , $j = 1, 2, \dots, n$, of random variables, Y_j , $j = 1, 2, \dots, n$, for the model (HL74.3).

3. The Set of All Possible Samples from a Population

Our interpretation of model quantities involves the idea of repeating over and over processes like selecting and measuring; for example, for the random variable Y_j , representing the response of the j th unit selected equiprobably for the sample, we say:

Y_j is a random variable whose distribution represents the possible values of the measured response variate for the j th unit in the sample of n units selected equiprobably from the respondent population, if the selecting and measuring processes were to be repeated over and over.

To pursue this idea of repeating over and over (or of *repetition*), we first discuss the set of all possible samples of size n that can be selected (without replacement) from a population of size N . There are $N^{(n)}$ such samples if the *order* of selection is taken into account; for example, there are $4^{(2)} = 12$ such samples of size 2 for a population of 4 elements.

Example HL74.1: A respondent population of $N = 4$ elements has the following integer *true* Y -values for its response variate: 1, 2, 4, 5 (so that: $\bar{Y} = 3$, $S = 1.8257$).

Table HL74.1 at the right lists the 12 possible ordered samples of size $n = 2$ that can be selected; *equiprobable* selecting (EPS) is assumed. For discussion below, the values of the average ($\bar{y}$) and sample error for each sample are also given.

Looking down the column of values for the *first* unit in the 12 samples, we see they are three copies of the values of the response variate of the elements of the respondent population; we see the same in the column of 12 values for the *second* unit. This result holds in general – in the set of $N^{(n)}$ possible ordered samples of size n , selected from a population of size N , the j th unit comprises $(N-1)^{(n-1)}$ copies of the values of the response variate of the elements of the respondent population. Hence, *under EPS*, the values of the unit in *any* position j of the set of all possible ordered samples of size n can be modelled like the values of the response variate of the elements of the respondent population – see equation (HL74.6) at the right; in statistics, it is more convenient to write equation (HL74.6) as the *response model* (HL74.7).

Table HL74.1

EPS of $n = 2$ units

Sample	$\bar{y}$	Error
(1, 2)	1½	-1½
(2, 1)	1½	-1½
(1, 4)	2½	-½
(4, 1)	2½	-½
(1, 5)	3	0
(5, 1)	3	0
(2, 4)	3	0
(4, 2)	3	0
(2, 5)	3½	½
(5, 2)	3½	½
(4, 5)	4½	1½
(5, 4)	4½	1½

$$Y_j \sim N(\mu, \sigma_s), \quad j = 1, 2, \dots, n, \quad \text{EPS} \quad \text{----(HL74.6)}$$

$$Y_j = \mu + \tau R_j, \quad j = 1, 2, \dots, n, \quad \tau R_j \sim N(0, \sigma_s), \quad \text{independent, EPS} \quad \text{----(HL74.7)}$$

NOTES: 1. The context of Example HL74.1 is *true* response variate values, indicated by the suffix τ on the random variables τY_j and τR_j (the **residuals**); conceptually, we are in the *first* sample ellipse in the schema at the lower right overleaf on page HL74.1 (see also Note 8 on page HL74.4).

- Another difference between equations (HL74.1) and (HL74.6) is the subscript s on σ_s , which indicates the variation quantified by this (probabilistic) standard deviation arises from the *selecting* (or ‘sampling’) process.
- 2. Because any sample position contains *multiple* copies (not just one copy) of the values of the respondent population responses, its (data) standard deviation is slightly *smaller* than the respondent population (data) standard deviation S , which is represented by σ_s in the model. This discrepancy is of no consequence for the present discussion in the usual practical situation of N large and $N \gg n$ (see also Note 4 on the facing page HL74.3).

3. The second and third columns of Table HL74.1 in Example HL74.1 above illustrate two other matters.

- Sample error – the difference between the value of the sample and respondent population attributes (here, averages) – has an *average of zero* over the set of all possible samples; stated another way, the average of the set of all possible sample averages is the respondent population average. This is the meaning of saying that, under EPS, the sample average is an *unbiased estimator* of the respondent population average. Symbolically, we write:

$$E(\bar{Y}) = \bar{Y} \quad \text{or} \quad E(\bar{Y}) - \bar{Y} = 0 \quad [\text{or, in the model: } E(\bar{Y}) = \mu \quad \text{or} \quad E(\bar{Y}) - \mu = 0]. \quad \text{----(HL74.8)}$$

This matter is illustrated more broadly in Appendix 3, which starts on page HL74.5 of this Highlight #74.

- EPS is needed for unbiasedness when estimating $\bar{Y}$ to ensure that all samples are *equally likely*, as reflected in calculating the *average* sample error.
- For the four Y -values of the response variate in this population, there *are* samples whose average is *equal* to the respondent population average, $\bar{Y}$; however, there are many populations for which *no* sample of a given size (or, conceivably, of *any* size other than N) has an average of $\bar{Y}$; i.e., *no* sample has *zero* sample error.
 - This is one reason a numerical Answer obtained by sampling is likely to be at least a little off the truth.

RESPONSE MODELS IN STAT 231: Modelling Sample Error and Measurement Error (continued 1)

4. A Probability Model for Selecting Equiprobably

Example HL74.1 on the facing page HL74.2 illustrates properties of the set of all possible samples of a population under EPS; we now argue from this case to the process of selecting equiprobably (a sample) *over and over*. Provided ‘over and over’ means selecting a number of times that is *large* in relation to $(N-1)^{(n-1)}$, because any sample is *equally* likely, the set of samples will be (roughly) many copies of the set of all possible samples. Subject to this new caveat of not having *exactly* equal proportions of each possible sample in the set, together with a slightly smaller standard deviation (see Note 2 on the facing page HL74.2), we can again use the model (HL74.7). Hence, for the behaviour of the sample *average* when selecting over and over, we use the model (HL74.3), rewritten above as equation (HL74.9).

- * The idea of obtaining, when selecting over and over, roughly *equal* proportions of each member of the set of all possible samples is like tossing a fair coin over and over, where we expect to observe roughly *equal* proportions of heads and tails.

NOTES: 4. Note 2 and the paragraph above refer to the standard deviation of τY_j as being slightly *smaller* than σ_s ; using a different approach in Statistical Highlight #77, the standard deviation of τY_j is actually as given in equation (HL74.10); the multiplier of $(1-1/N)$ under a square root is obviously close to 1 in value for most populations encountered in practice, which have *many* elements.

$$\tau \bar{Y} \sim N(\mu, \sigma_s \sqrt{\frac{1}{n}}), \text{ EPS} \quad \text{-----(HL74.9)}$$

$$s.d.(\tau Y_j) = S \sqrt{1 - \frac{1}{N}} \quad \text{-----(HL74.10)}$$

$$s.d.(\tau \bar{Y}) = S \sqrt{\frac{1}{n} - \frac{1}{N}} \quad \text{-----(HL74.11)}$$

- The consequence of equation (HL74.10) is that the standard deviation of $\tau \bar{Y}$ is as given in equation (HL74.11) above; in the usual situation in practice, where $N \gg n$, this standard deviation becomes effectively that of equation (HL74.9), where σ_s is the model parameter representing the respondent population (data) standard deviation S .
- 5. The assumption of a *normal* model for the shape of the histogram of the (true) response variate Y -values in the respondent population can be relaxed if we can assume that the Central Limit Theorem will provide adequate *approximate* normality of the random variable $\tau \bar{Y}$ representing the sample average $\tau \bar{Y}$ under repetition.

5. An EPS-Based Probability Model for Measuring

The response model (HL74.7), with *normal* residuals, for equiprobable selecting over and over can be adapted to model the process of *measuring* (independently) over and over the value of a variate for a unit, for two reasons:

- * We know from histograms like those for the paper thickness data in Statistical Highlight #39 (and for the coin weights in Figure 6.2 of the STAT 231 2004 Course Materials) that normal distributions are a reasonable model for values produced by (some) measuring processes.
- * We can model measuring as a selecting process: possible values for the variate being measured could be written on slips of paper, the slips placed in a box and a slip selected by EPS; the number on the slip is then regarded as the ‘measured’ value; measuring over and over would be modelled as selecting slips equiprobably *with* replacement over and over.

We therefore write the response model for measuring as equation (HL74.12) at the right, where:

$${}_MY = \tau + {}_MR, \quad {}_MR \sim N(0, \sigma_M), \text{ independent, EPS} \quad \text{-----(HL74.12)}$$

${}_MY$ is a random variable whose distribution represents the possible values of the measurement of the response variate of a unit, if the measuring process were to be repeated over and over on this unit.

τ is a model parameter which represents the *true value* of the response variate of the unit being measured.

δ is a model parameter (called the *bias*) which represents the *inaccuracy* of the measuring process; the value of δ *quantifies* the inaccuracy of the measuring process – as inaccuracy *increases* (i.e., as accuracy *decreases*), δ *increases*.

R is a random variable (called the *residual*) whose distribution represents the possible *differences*, from the structural component of the model, of the value of the measurement of the response variate of the unit being measured, if the measuring process were to be repeated over and over on this unit.

σ_M the *standard deviation* of the normal model for the distribution of the residual, is a model parameter (called the *variability*) which represents the *imprecision* of the measuring process and describes measuring variation if the measuring process were to be repeated over and over on a unit; the value of σ_M *quantifies* the imprecision of the measuring process – as imprecision *increases* (i.e., as precision *decreases*), σ_M *increases*.

For a *calibrated* measuring system, where the inaccuracy has been (essentially) removed using a defined standard, the parameter δ can be omitted and the model written as equation (HL74.13).

$${}_MY = \tau + {}_MR, \quad {}_MR \sim N(0, \sigma_M), \text{ independent, EPS} \quad \text{-----(HL74.13)}$$

NOTES: 6. The absence of the subscript j on ${}_MY$ and ${}_MR$ in the response models (HL74.12) and (HL74.13) is because we are modelling the process of measuring *once* the value of a variate for a unit.

- 7. The suffix M on ${}_MY$ (and ${}_MR$) in the models (HL74.12) and (HL74.13) reminds us we are dealing with a *measured* value of a response variate, as distinct from τY_j in equation (HL74.7), the corresponding *true* value for the j th unit in the sample.

6. A Probability Model for Selecting Equiprobably and Measuring

To obtain a response model for both selecting *and* measuring, two processes responsible for *incomplete* data and, hence, for *uncertain* Answers, we combine the response models (HL74.7) and (HL74.13); the resulting model [equation (HL74.16) below] describes the processes of EPS of n units from a respondent population (for which $N \gg n$) and measuring their response variate values once each with a *calibrated* measuring system.

Equation (HL74.7) is given again at the right, and equation (HL74.13) is rewritten as equation (HL74.14), where τ has been replaced by the *random variable* τY_j because we are measuring (once) the response of the j th unit in the sample *selected by EPS*; ${}_MY_j$ has a subscript j for the same reason. Also, the suffixes τ and M on the residuals are to distinguish the residuals in the two models for selecting and measuring; these two sets of residuals become *one* set denoted R_j in the combined model (HL74.16) – see equation (HL74.17) below.

Then, writing ${}_MY_j$ in equation (HL74.14) as its expression in (HL74.7), we obtain equation (HL74.15), which we rewrite as equation (HL74.16), where σ quantifies the *overall* variation due to *both* selecting and measuring [equation (HL74.17)].

$${}_TY_j = \mu + {}_TR_j, \quad j=1, 2, \dots, n, \quad {}_TR_j \sim N(0, \sigma_s), \quad \text{independent, EPS} \quad \text{----(HL74.7)}$$

$${}_MY_j = {}_TY_j + {}_MR_j, \quad j=1, 2, \dots, n, \quad {}_MR_j \sim N(0, \sigma_m), \quad \text{independent, EPS} \quad \text{----(HL74.14)}$$

$${}_MY_j = \mu + {}_TR_j + {}_MR_j, \quad j=1, 2, \dots, n \quad \text{----(HL74.15)}$$

$${}_MY_j = \mu + R_j, \quad j=1, 2, \dots, n, \quad R_j \sim N(0, \sigma), \quad \text{independent, EPS} \quad \text{----(HL74.16)}$$

$$R_j = {}_TR_j + {}_MR_j, \quad \sigma = \sqrt{\sigma_s^2 + \sigma_m^2} \quad \text{----(HL74.17)}$$

NOTES: 8. In this Highlight #74, we have developed a model for selecting *and* measuring by combining separate models for the two processes; in doing so, we distinguish true and measured response variate values (${}_TY_j$ and ${}_MY_j$). After discussing this distinction here, the data we encounter elsewhere are essentially *always* obtained by measuring and explicit consideration of the true value ${}_TY_j$ of equation (HL74.7) seldom arises. For convenience, we will therefore usually write the model (HL74.16) as at the right without the suffix M .

$$Y_j = \mu + R_j, \quad j=1, 2, \dots, n, \quad R_j \sim N(0, \sigma), \quad \text{independent, EPS} \quad \text{----(HL74.16)}$$

- ${}_MY_j$ and Y_j in equation (HL74.16) is the quantity defined early in Section 3 on the second side (page HL74.2) of this Highlight #74 – this definition involves *both* selecting and measuring; conceptually, we have reached the *second* sample ellipse in the schema at the lower right of the first side (page HL74.1) of this Highlight #74.

9. Based on the model (HL74.16), equation (HL74.18) at the right describes the behaviour of the random variable representing the sample average, $\bar{Y}$, subject to variation from *both* equiprobable selecting and measuring.

$$\bar{Y} \sim N(\mu, \sigma/\sqrt{n}), \quad \text{EPS} \quad \text{----(HL74.18)}$$

- Despite the apparent similarity of equations (HL74.3) and (HL74.18), their derivations and the interpretation of $\bar{Y}$, μ and σ differ in statistically vital ways.

7. Appendix 1: A Comparison of Approaches to Modelling the Behaviour of $\bar{Y}$

Three approaches to modelling the behaviour of the sample average, as a basis for estimating the respondent population average, are discussed in this Highlight #74; this Appendix 1 presents an overview of the strengths and weaknesses of each.

- * *Probability modelling* starts with the normal model (HL74.1) and then uses probability theory for random variables to obtain the model (HL74.3); the weakness of this approach for *statistics* is that the ideas of Sections 2 to 6, and the assumptions they make explicit, are only *implicit*. It is likely to be unclear to the beginning student why statements like equations (HL74.1) and (HL74.3) can be used to describe the real-world processes of selecting (and measuring, if it is even considered in this approach) and it is easy for statistical practice based on the probability approach to overlook essential components of the Plan.

$$Y \sim N(\mu, \sigma) \quad \text{----(HL74.1)}$$

$$\bar{Y} \sim N(\mu, \sigma/\sqrt{n}) \quad \text{----(HL74.3)}$$

- * The *response model* (HL74.16), as developed on the basis of the *same* probability models (HL74.1) and (HL74.3), recognizes explicitly Question formulation (via consideration of the target, study and respondent populations) and properties of the selecting and measuring processes needed as a basis for the relevant probability models, written as the response models (HL74.7) and (HL74.14). This approach therefore overcomes the weaknesses of the probability modelling approach, albeit at the cost of more detail in presentation; the final probability expression for the behaviour of $\bar{Y}$ involving variation from *both* selecting and measuring, is equation (HL74.18). A weakness of this approach is that it does not recognize explicitly the finite size (N elements) of the respondent population.

$${}_TY_j = \mu + {}_TR_j, \quad j=1, 2, \dots, n, \quad {}_TR_j \sim N(0, \sigma_s), \quad \text{indep., EPS} \quad \text{----(HL74.7)}$$

$${}_MY_j = {}_TY_j + {}_MR_j, \quad j=1, 2, \dots, n, \quad {}_MR_j \sim N(0, \sigma_m), \quad \text{indep., EPS} \quad \text{----(HL74.14)}$$

$$Y_j = \mu + R_j, \quad j=1, 2, \dots, n, \quad R_j \sim N(0, \sigma), \quad \text{indep., EPS} \quad \text{----(HL74.16)}$$

$$\bar{Y} \sim N(\mu, \sigma/\sqrt{n}), \quad \text{EPS} \quad \text{----(HL74.18)}$$

- * The so-called *design-based* approach developed in Statistical Highlight #77, based on the equiprobable selecting of the sampling protocol, incorporates the finite size (N elements) of the respondent population, but it cannot easily take direct account of measurement error. Its expression for the behaviour of $\tau\bar{Y}$ is equation (HL74.19); the *approximate* normality comes from the Central Limit Theorem approximation.

$$\tau\bar{Y} \div N(\bar{Y}, \mathbf{S}/\sqrt{1 - \frac{1}{N}}), \quad \text{EPS} \quad \text{----(HL74.19)}$$

Estimating $\bar{Y}$ using the models (HL74.3), (HL74.18) and (HL74.19) requires an *estimate* of σ or $\mathbf{S}$, customarily taken as the *sample*

RESPONSE MODELS IN STAT 231: Modelling Sample Error and Measurement Error (continued 2)

standard deviation, s . Because s is calculated from *measured* values of the response variate, its value *includes* variation due to measuring *as well as* selecting; this inclusion is modelled in (HL74.18) but is a *fortuitous* benefit for (HL74.3) and (HL74.19).

Thus, the model (HL74.18), derived from the response model (HL74.16), provides a basis for statistical practice that is missing from (HL74.3); the additional insight from (HL74.19) is that the multiplier of $\mathbf{S}$ in the standard deviation of $\tau\bar{Y}$ has weak dependence on $\mathbf{N}$; ignoring $\mathbf{N}$ in (HL74.18) is justified provided the sample size is a *small proportion* of the respondent population size ($n \ll \mathbf{N}$).

8. Appendix 2: Using the Response Model for Measuring

When investigating the inaccuracy and imprecision of a measuring system [often as part of broader data-based investigating to answer Question(s) of interest], a common approach is to make m independent measurements of the same value. The appropriate response models are shown at the right as equations (HL74.20) and (HL74.21); the former is applicable when measuring a *known* value to investigate both inaccuracy *and* imprecision, whereas the latter is used when investigating only imprecision; the two forms of (HL74.20) differ in whether the measured value or its difference from the true value is taken as the response variate – one or the other may be more convenient, depending on context.

$${}_mY_j = \tau + \delta + {}_mR_j, \quad j = 1, 2, \dots, m, \quad {}_mR_j \sim N(0, \sigma_m), \quad \text{independent, EPS} \quad \text{---(HL74.20)}$$

$${}_mY_j - \tau = \delta + {}_mR_j, \quad j = 1, 2, \dots, m, \quad {}_mR_j \sim N(0, \sigma_m), \quad \text{independent, EPS} \quad \text{---(HL74.20)}$$

$${}_mY_j = \tau + {}_mR_j, \quad j = 1, 2, \dots, m, \quad {}_mR_j \sim N(0, \sigma_m), \quad \text{independent, EPS} \quad \text{---(HL74.21)}$$

Example HL74.2: A medical laboratory purchases a standard, certified to contain 200 mg/dL of cholesterol, to calibrate its process for measuring serum cholesterol levels; data from 15 measurements taken over an 8-hour day were:

Table	Measured value (${}_mY_j$)	194	196	202	198	195	188	203	200	196	195	198	202	203	201	199	Av.	S.d.
HL74.2:	Measured – true value (${}_mY_j - \tau$)	-6	-4	2	-2	-5	-12	3	0	-4	-5	-2	2	3	1	-1	-2	4.1231

These data show that: the estimate of measuring *inaccuracy* (represented in the model by δ) is -2 mg/dL;
the estimate of measuring *imprecision* (represented in the model by σ_m) is 4.1 mg/dL.

With inaccuracy estimated to be about half the magnitude of imprecision, the former may have little practical importance here, but more detailed interpretation of these data requires extra-statistical knowledge.

NOTE: 10. Under the idealizations of the models (HL74.20) and (HL74.21), it can be shown that, for measuring processes:

- inaccuracy is *unaffected* by averaging – looking ahead, one average is the estimate of the *intercept* of the centred form of the straight-line model in simple linear regression (see page HL73.2 in Statistical Highlight #73).
- inaccuracy is *removed* by differencing – differences occur when comparing and when calculating (data) standard deviations and the *slope* of the straight-line regression model (see page HL73.2 in Statistical Highlight #73).

9. Appendix 3: The Set of All Possible Samples from a Population

Section 3 introduced the idea of all possible samples that can be selected from a population. The sets of ten and nine histograms on pages HL74.6 and HL74.7 illustrate this idea for a population of $\mathbf{N} = 10$ elements with response variate values $\mathbf{Y} = 1, 2, \dots, 10$. These histograms, for sample averages and sample (data) standard deviations, show important properties of the selecting process.

- * The sample average takes a value which is determined by the particular sample selected; these values, over the set of all possible samples of a given size, form a *distribution*; under EPS, *this is the distribution of the random variable $\bar{Y}$* , a result that could be called **The Fundamental Theorem of Statistics** by analogy with The Fundamental Theorem of Calculus.
 - The mean of this distribution is the *population average* – this is the *unbiasedness* under EPS of the sample average as an estimator of the population average (see the second column of Table HL74.3 in Appendix 4 on page HL74.8).
 - The standard deviation of this distribution is given by equation (HL74.11) except that, because $s.d.(\bar{Y})$ is a *probabilistic* standard deviation and $s_{\bar{Y}}$ (given in the third column of Table HL74.3 in Appendix 4) is a *data* standard deviation, the former is $\sqrt{k/(k-1)}$ times the latter, where k is $\binom{\mathbf{N}}{n}$, the number of possible (unordered) samples of size n .
- * The distribution of sample averages reminds us that for the *particular* sample selected when executing the Plan for an investigation, its sample error is *unknown* – its attribute may be close to the population attribute or (sometimes) it may not be; that is, an Answer has sampling *uncertainty*, meaning we cannot *know* how close an Answer obtained by sampling is to the truth.
- * The *decreasing variation* of the values of sample averages with *increasing* sample size, visible as the decreasing width down the page of the histograms of sample averages, illustrates sampling *imprecision* decreasing with increasing *sample size*; *i.e.*, under EPS, increasing sample size decreases the average magnitude of sample error and so decreases sampling uncertainty.
- * The change in distribution *shape*, most noticeable as the sample size increases from 1 to 2, illustrates the *idea* behind the Central Limit Theorem – a *non-normal* (uniform) distribution starting to show the central ‘heaping’ of the normal distribution.
 - The heaping, which first appears when $n = 2$, persists up to $n = 8$ and then disappears; this is likely the effect of *dependence* among samples, which consist of most of the elements in the population, becoming dominant over the heaping process reminiscent of the Central Limit Theorem.

(continued overleaf)

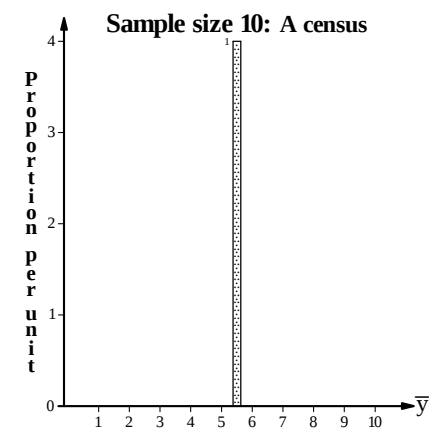
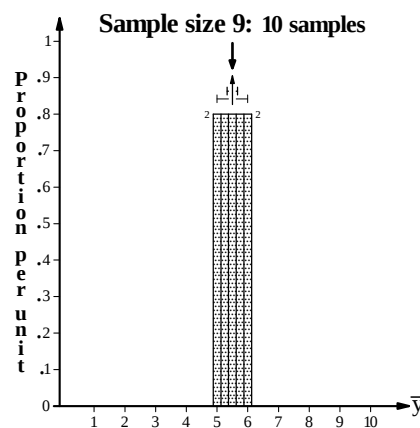
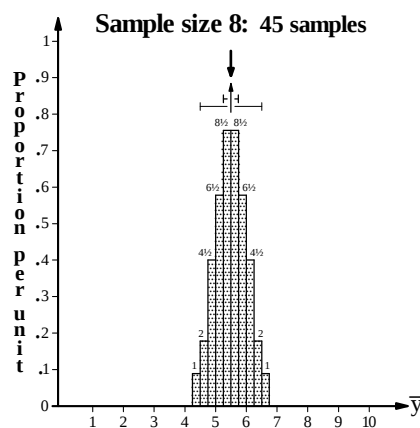
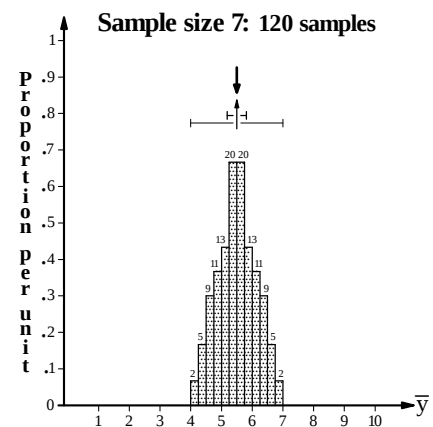
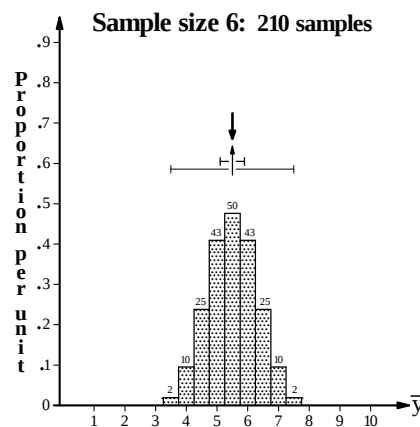
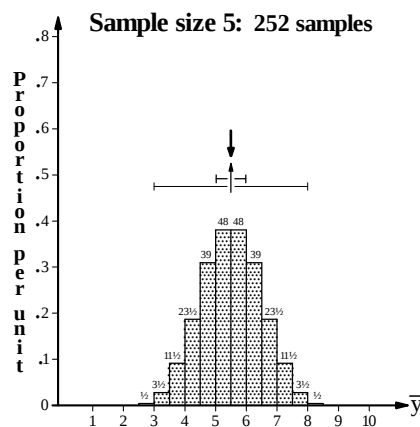
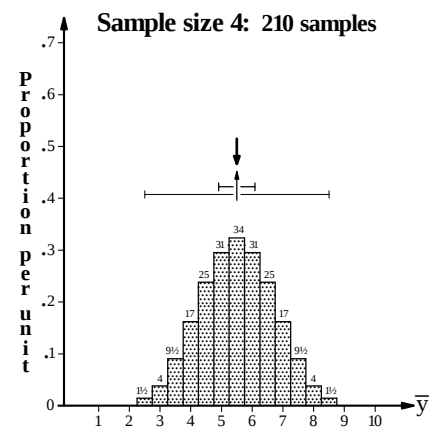
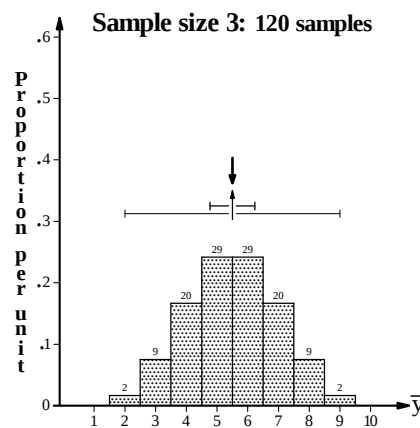
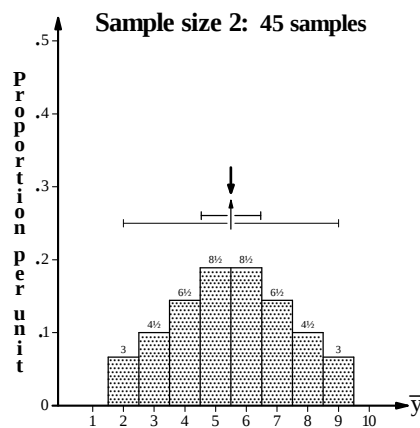
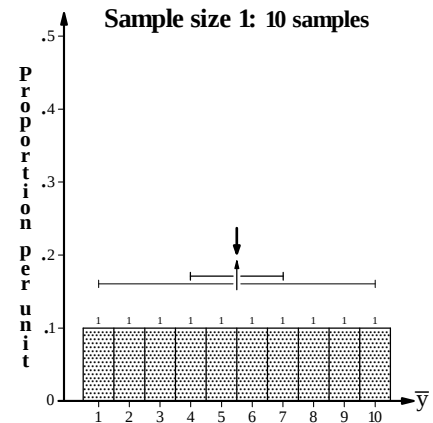
9. Appendix 3: The Set of All Possible Samples from a Population (continued)

These ten histograms show the distribution of the *averages* of all possible samples of a given size that can be selected without replacement from a population of $N = 10$ elements, whose response variate values are $Y = 1, 2, \dots, 10$. [A histogram of these ten population *responses* would be like that given at the right for the averages of the ten possible samples of size 1.]

The decreasing vertical axis scale unit down the page reduces the visual impact of increasing histogram height; bar *frequencies* are given by the numbers at the tops of the bars.

Above each histogram, except the one at the bottom right of the page:

- the upper bold arrow ($\downarrow$) indicates the value of the *population average* $\bar{Y} = 5.5$;
- the lower arrow ($\uparrow$) indicates the *average* of the set of *sample averages* ($\bar{\bar{y}} = 5.5$);
- the upper bar (—) crossing the lower arrow has a length equal to the value of the (data) *standard deviation* of the set of sample averages;
- the lower such bar shows the *range* of the set of sample averages.



RESPONSE MODELS IN STAT 231: Modelling Sample Error and Measurement Error (continued 3)

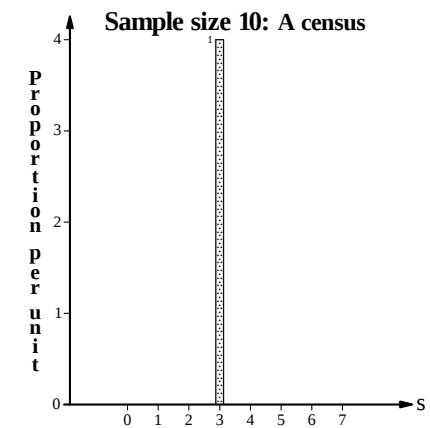
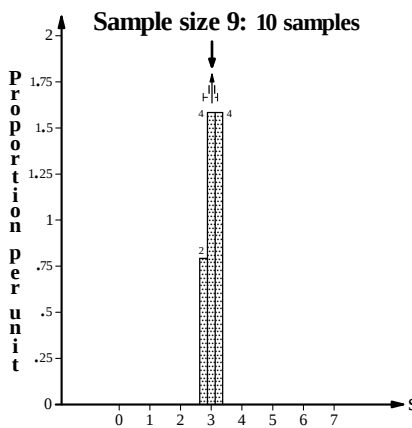
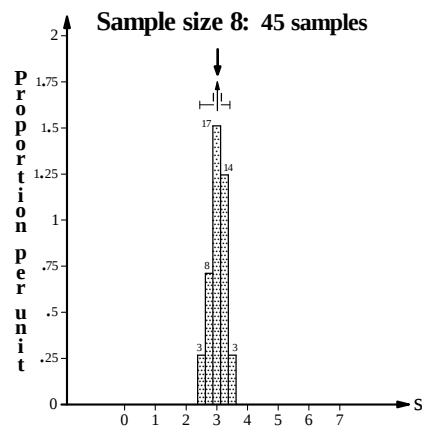
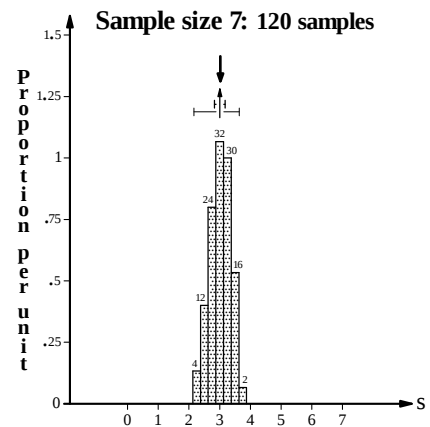
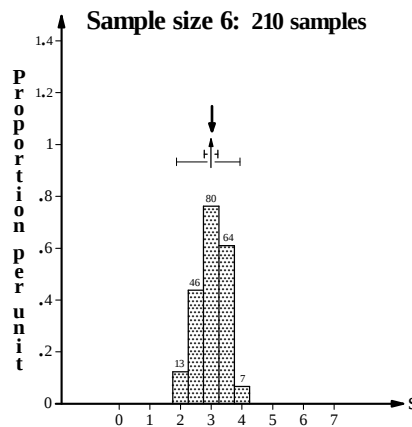
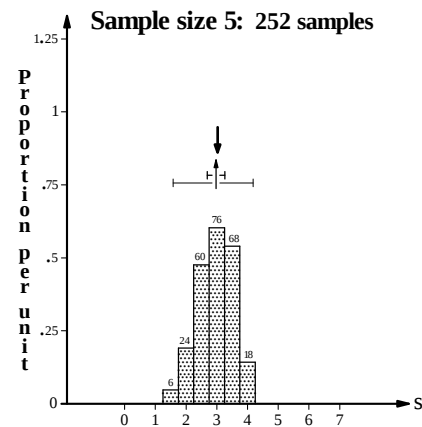
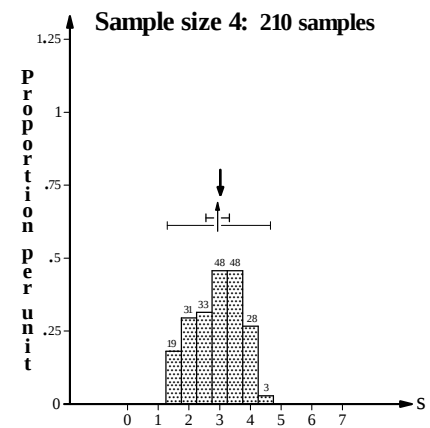
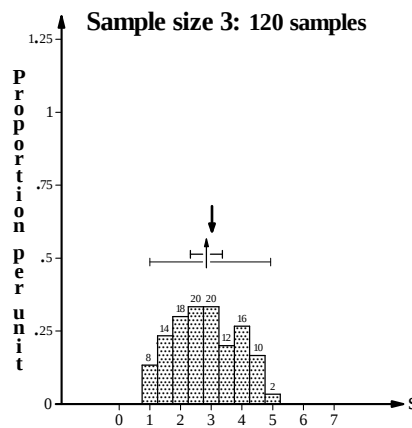
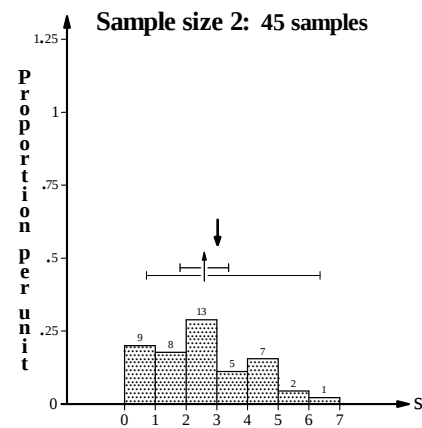
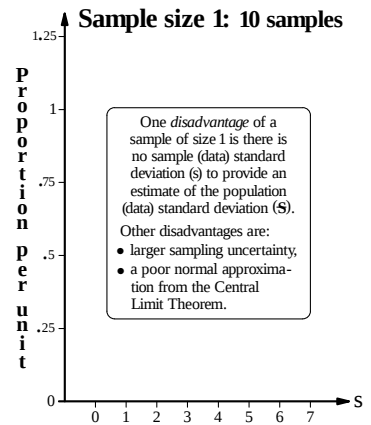
These nine histograms show the distribution of the (data) *standard deviations* of all possible samples of a given size that can be selected without replacement from a population of $N=10$ elements, whose response variate values are $Y=1, 2, \dots, 10$.

[A histogram of these ten population *responses* would be like that given at the top right of the facing page HL74.6 for the averages of the ten possible samples of size 1.]

The decreasing vertical axis scale unit down the page reduces the visual impact of increasing histogram height; bar *frequencies* are given by the numbers at the tops of the bars.

Above each histogram, except the one at the bottom right of the page:

- the upper bold arrow ($\downarrow$) indicates the value of the *population* standard deviation $S = 3.027\ 650$;
- the lower arrow ($\uparrow$) indicates the *average* ($\bar{s}$) of the set of *sample* (data) standard deviations;
- the upper bar (—) crossing the lower arrow has a length equal to the value of the (data) *standard deviation* of the set of sample (data) standard deviations;
- the lower such bar shows the *range* of the set of sample (data) standard deviations.



9. Appendix 3: The Set of All Possible Samples from a Population (continued)

- * The histograms of sample (data) standard deviations illustrate two differences from sample averages:
- their distributions are narrower (even with a *smaller* scale unit on the horizontal axis) and are not symmetrical;
 - the sample (data) standard deviation is *not* an unbiased estimator of the population (data) standard deviation, but the estimating bias decreases as sample size increases [see Appendix 3 on pages HL77.9 and HL77.10 in Statistical Highlight #77.]
 - Estimating bias is *unlike* sampling and measuring inaccuracy which do *not* decrease with increasing sample size.

10. Appendix 4: Data Sets for Appendix 3

The histograms in Appendix 3 were constructed from the data tabulated below and on the facing page HL74.9. The average ($\bar{y}$) and (data) standard deviation ($s_{\bar{y}}$) of the sample averages for each sample size (n) are given in the second and third columns of the table below, as are the average ($\bar{S}$) and (data) standard deviation (s_s) of the sample (data) standard deviations on the facing page HL74.9.

Table HL74.3
Ordered Sample Averages ($\bar{y}$)

n	$\bar{y}$	$s_{\bar{y}}$	Ordered Sample Averages ($\bar{y}$)																																																																																																																																																																																																																																																																																							
1	11/2	3.0277	1	2	3	4	5	6	7	8	9	10																																																																																																																																																																																																																																																																														
2	11/2	1.9365	3/2 6	2 6	5/2 6	5/2 6	3 13/2	3 13/2	7/2 13/2	7/2 13/2	7/2 7	4 7	4 7	4 15/2	9/2 15/2	9/2 15/2	9/2 8	9/2 8	5 17/2	5 17/2	5 9	5 19/2	11/2	11/2	11/2	11/2	11/2																																																																																																																																																																																																																																																															
3	11/2	1.4686	2 13/3 16/3 6 7	7/3 13/3 16/3 6 7	8/3 13/3 16/3 6 7	8/3 13/3 16/3 6 7	3 13/3 16/3 19/3 22/3	3 13/3 16/3 19/3 22/3	10/3 14/3 16/3 19/3 22/3	10/3 14/3 16/3 19/3 22/3	10/3 14/3 16/3 19/3 22/3	10/3 14/3 16/3 19/3 22/3	10/3 14/3 16/3 19/3 22/3	11/3 14/3 17/3 19/3 23/3	11/3 14/3 17/3 19/3 23/3	11/3 14/3 17/3 19/3 23/3	11/3 14/3 17/3 19/3 23/3	11/3 14/3 17/3 19/3 23/3	11/3 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5 17/3 20/3 25/3	4 5

ACKNOWLEDGEMENTS: The numbers in the tables in Appendix 4 were kindly calculated by Professor R.W. Oldford using Quail. The development of the model for sample and measurement error in Section 6 on page HL74.4 is based on ideas in Chapter 5 of the 2004 Course Notes for Statistics 231, by Professors R.W. Oldford and R.J. MacKay.

#HL74.9

11. Appendix 5: Sampling Protocols Beyond EPS [The title matter of Statistical Highlight #85.]

Equiprobable (or simple random) selecting of units consisting of *individual elements* from an *unstratified* (respondent) population is useful for modelling the selecting process but, in practice, more complex sampling protocols are used. Two such protocols are:

- * **cluster selecting:** selecting equiprobably units from the (respondent) population that are *groups* of elements – clusters may be of *equal* size (e.g., cardboard boxes of 24 cans of soup) or *unequal* size (e.g., households);
- * **stratified selecting:** subdividing the (respondent) population into groups (called **strata**) so that elements *within* a stratum have *similar* response variate values ('homogeneity of strata') and elements in different strata *differ* as much as practicable from each other; the sample is obtained by equiprobable selecting of units consisting of individual elements from *each* stratum.

[The element-unit distinction is discussed in Appendix 1 on pages HL77.8 and HL77.9 in Statistical Highlight #77.]

Example HL74.3 below illustrates the effects of *clustering* and *stratifying* on *sampling imprecision*, also bearing in mind that:

- *clustering* is commonly used because of the availability of a clustered *frame*, thereby avoiding the cost of generating a frame as part of the investigating – a **frame** can be thought of as a *list* of the (respondent) population sampling units (here, clusters);
- *stratifying* is commonly used because it provides (the often useful additional) *subdivision* of Answers by stratum.

Example HL74.3: A respondent population of $N = 4$ elements (or units) has the following integer $\mathbf{Y}$ -values for its response variate:

1, 2, 4, 5 [so that the population average and (data) standard deviation are: $\bar{Y} = 3$, $S \approx 1.8257$];

we examine the *sampling imprecision*, under equiprobable selecting (EPS) with a sample size of $n = 2$, of the random variable $\bar{Y}$ whose values are the sample average $\bar{y}$, as an estimator of $\bar{Y}$, using three sampling protocols:

- EPS of two units, each consisting of one element, from the *unstratified* population;
- EPS of one *cluster*, of size $L = 2$ elements, from the *unstratified* population;
- EPS of one unit, consisting of one element, from each of two *strata* of size $N_1 = N_2 = 2$.

Note that *each* estimator is *unbiased*, because $E(\bar{Y}) = \bar{Y}$ or $E(\bar{Y}) - \bar{Y} = 0$ [the *average* sample error is zero].

Table HL74.5

Unstratified population
EPS of two elements

Sample	$\bar{y}$	Error
(1, 2)	1½	–1½ large
(1, 4)	2½	–½ medium
(1, 5)	3	0 small
(2, 4)	3	0 small
(2, 5)	3½	½ medium
(4, 5)	4½	1½ large

Designation of sample error as *large*, *medium* or *small* is only for convenience in the context of Example HL74.3.

Table HL74.6a

Unstratified population
Clusters: [1, 2], [4, 5]

EPS of one cluster (L = 2)

Sample	$\bar{y}$	Error
(1, 2)	1½	–1½ large
(4, 5)	4½	1½ large

Table HL74.6b

Unstratified population
Clusters: [1, 4], [2, 5]

EPS of one cluster (L = 2)

Sample	$\bar{y}$	Error
(1, 4)	2½	–½ medium
(2, 5)	3½	½ medium

Table HL74.6c

Unstratified population
Clusters: [1, 5], [2, 4]

EPS of one cluster (L = 2)

Sample	$\bar{y}$	Error
(1, 5)	3	0 small
(2, 4)	3	0 small

Table HL74.7a

Stratified population
Strata: [1, 2], [4, 5]

EPS of one element per stratum

Sample	$\bar{y}$	Error
(1, 4)	2½	–½ medium
(1, 5)	3	0 small
(2, 4)	3	0 small
(2, 5)	3½	½ medium

Table HL74.7b

Stratified population
Strata: [1, 4], [2, 5]

EPS of one element per stratum

Sample	$\bar{y}$	Error
(1, 2)	1½	–1½ large
(1, 5)	3	0 small
(2, 4)	3	0 small
(4, 5)	4½	1½ large

Table HL74.7c

Stratified population
Strata: [1, 5], [2, 4]

EPS of one element per stratum

Sample	$\bar{y}$	Error
(1, 2)	1½	–1½ large
(1, 4)	2½	–½ medium
(2, 5)	3½	½ medium
(4, 5)	4½	1½ large

Example HL74.3 illustrates that:

- The effect on (sampling) imprecision of clustering and of stratifying depends on how each is *implemented* in the Plan – that is, it depends on this component of the *sampling protocol*.
- Clustering and stratifying affect imprecision by determining which of the *possible* samples of size n have non-zero selecting probabilities.
- Decreased imprecision is favoured by *heterogeneity of clusters* but by *homogeneity of strata* with respect to the response(s) of interest.
 - In the middle and right-hand columns of Example HL74.3, heterogeneity increases *down* the three clustered sampling protocols, homogeneity increases *up* the three stratified protocols.

[As an exercise, quantify the sample error *variation* by calculating the relevant (data) *standard deviation* for each of the seven sampling protocols; comment on what is illustrated by the values obtained.]

- There is a sense in which clustering is *passively* accepted in the interests of reducing investigation cost, whereas stratifying may be *actively* imposed by the investigator(s) on [or may be a natural feature of] the study (or respondent) population.
- While EPS from an *unstratified* population implies *equal* inclusion probabilities for all population *elements*, the converse does *not* hold – in the three clustered and three stratified sampling protocols, all *elements* have equal inclusion probabilities but all six *samples* of size 2 are *not* equally probable [four samples and two samples (respectively) have *zero* selecting probability].