

Chapter 2

Types of Errors in Instrumental Analysis

2.1 Overview

Even under constant experimental conditions (same operator, same tools, and same laboratory, short time intervals between the measurements), repeated measurements of series of identical samples always lead to results which differ among themselves and from the true value of the sample. Therefore, quantitative measurements cannot be reproduced with absolute reliability.

According to their character and magnitude, the following types of deviation can be distinguished [1–4]:

Random Errors. Random errors are the components of measurement errors that vary in an unpredictable manner in replicated measurements. They reflect the distribution of the results around the mean value of the sample which are randomly distributed to lower and higher values. Random errors characterize the *reproducibility* of measurements, and, therefore, their *precision*. They are caused by effects such as measuring techniques (e.g. noise), sample properties (e.g. inhomogeneities), and chemical effects (e.g. equilibrium). Even under carefully controlled conditions random errors cannot, in principle, be avoided, they can only be minimized and evaluated with statistical methods.

Systematic Errors. Systematic deviations (errors) displace the results of analytical measurements to one side, to higher or lower values which lead to false results. Such an effect is described by the performance characteristic *trueness*, which is defined as “the closeness of agreement between the expectation of a test result or measurement result and a true value” [1]. Measurement trueness is not a *quantity* and cannot be expressed numerically [2], but measures for closeness of agreement can be given. Thus the trueness can be quantified as *bias* which is defined as the difference between the average of several measurements on the same sample $\hat{x}$ and its (conventionally) true value μ :

$$\text{Bias}(x) = \hat{x} - \mu \quad (2.1-1)$$

or if expressed as a percentage

$$\text{Bias } \% = \frac{(\hat{\bar{x}} - \mu)}{\mu} \cdot 100 \quad (2.1-2)$$

or as the recovery ratio

$$\text{Rr}\% = \frac{\hat{\bar{x}}}{\mu} \cdot 100. \quad (2.1-3)$$

In contrast to random errors, systematic errors can and *must* be avoided or eliminated if their origins become known, because they yield false results. Note that systematic errors cannot be statistically evaluated.

Systematic errors are always combined with random errors as shown in Fig. 2.1-1.

The measurement *accuracy* is defined as the “closeness of agreement between a measured quantity value and a true quantity value of a measurand” [2]. The measurement accuracy is not given a numerical value, but it is a qualitative performance characteristic which expressed the closeness of agreement between a measurement result and the value of the measurand, and thus it describes the precision as well as the trueness [5]. Therefore, the term “measurement accuracy” should not be used for measurement trueness.

The performance parameter of accuracy is the measurement uncertainty.

Measurement Uncertainty. The uncertainty of measurements is defined as “a parameter associated with the result of a measurement that characterizes the dispersion of the value that could reasonably be attributed to the measurand” [2]. The uncertainty concept divides the errors into two uncertainty components:

- Those that can be characterized by the experimental *standard deviations* (uncertainty components from Type A).
- Those that can be evaluated from *assumed probability distributions* based on experimental or other information (uncertainty components from Type B).

The combined uncertainty from both components is calculated by the law of propagation of errors (see Chap. 10).

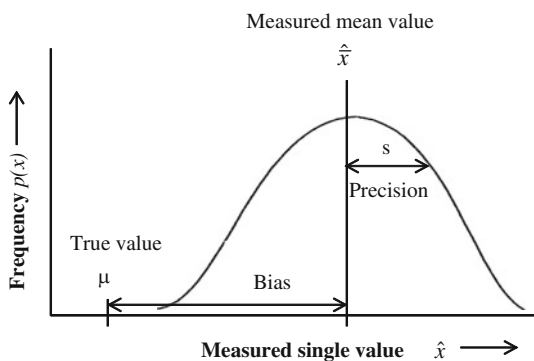


Fig. 2.1-1 Graphical representation of the terms precision and trueness

Outliers. Outliers are individual measurement values which considerably differ from the mean value. Outliers would falsify the estimation of parameters such as the mean value and the standard deviation, and therefore they must be detected by statistical methods and eliminated from the data set or, if this is not possible, one must work with methods resistant to outliers (robust methods [6]).

Trend. A data set shows a trend when the chronologically ordered values move steadily downwards or upwards. Such a data set is not under statistical control; therefore, after it has been recognized statistically, the trend must be eliminated. Note that a data set which shows a trend is to be rejected.

Gross Errors. Gross errors result from human mistakes, or have their origins in instrumental or computational errors. Frequently, they are easy to recognize and the origins must be eliminated.

Challenge 2.1-1

Table 2.1-1 shows series of data sets obtained by the five methods A–E. Which kinds of errors can be *visually* detected? The true content of the sample is $\mu = 100$.

Table 2.1-1 Hypothetical analytical results obtained with five methods

Method	x_1	x_2	x_3	x_4	x_5	x_6	$\bar{x}$
A	99	101	98	100	100	102	100
B	106	98	104	95	94	103	100
C	106	102	102	99	98	93	100
D	99	101	98	82	100	102	97
E	115	117	116	116	112	114	115

Solution to the Challenge 2.1-1

The data set in series C obviously shows a trend downwards, i.e. a trend is present. Though the calculated mean value is correct the data are not appropriate for the analytical result. The data set in C has to be rejected.

The value 82 in series D is obviously an outlier which leads to a false mean value. After elimination of this value a correct value (100) can be calculated.

Method E clearly yields a false mean value $\bar{x} = 115$. The result is obviously too high because a systematic error is present.

Methods A and B yield correct mean values but the individual results show a higher dispersion around that mean in series B than in A. This means that the precision in series A is better.

This exercise can be regarded as a *plausibility control*, which is an import step in analytical quality assurance. Plausibility control means checking data series,

analytical results, and others, without statistical tests, to see if the data *can be corrected*. This procedure has to be carried out before the release of data for further processes or for the documentation of analytical results. Thus, for example, the check of series D reveals the existence of an outlier ($x_4 = 82$) for which an outlier test must be carried out (see Challenge 3.3-1), and the trend in series C is also obvious.

Note that errors cannot always be clearly recognized; statistical methods are mostly necessary, but this is the subject of the following chapters.

2.2 Random Errors

2.2.1 Distribution of Measured Values

When one wants to view the distribution of many available data, it is useful to group the n data into k classes with n_j variables in each class and visualize their frequency density or probability distribution $p(x)$ with a *histogram*, which is a graphical display of tabulated frequencies presented as bars [7–9]. Figure 2.2.1-1 shows an example for the frequency density $p(x)$ of measured values x . The bars must be adjacent and the intervals (or bands) are generally of the same size. The rule of thumb $k = \sqrt{n}$ provides an appropriate number of classes k for the construction of a histogram with n data.

If the number of repeated measurements is increased to infinity and one reduces the width of the classes towards zero, a symmetrical bell-shaped distribution of measurement values is usually obtained, which is called *Gaussian* or *normal distribution* (see curve ND in Fig. 2.2.1-2).

The frequency density $p(x)$ is described by the function

$$p(x) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left\{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right\}. \quad (2.2.1-1)$$

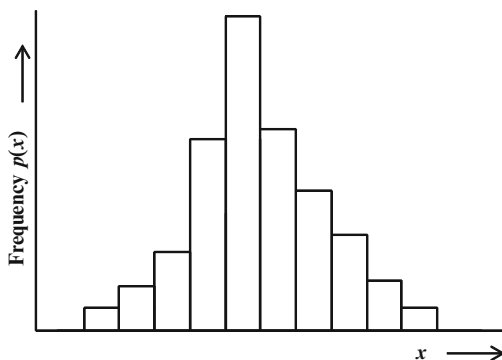


Fig. 2.2.1-1 An example for the frequency density $p(x)$ of measured values x

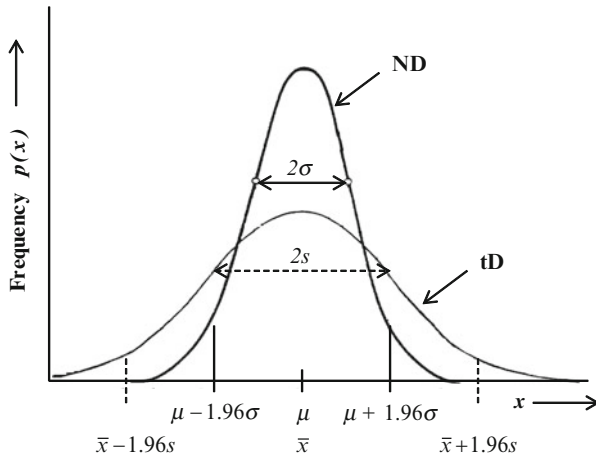


Fig. 2.2.1-2 Gaussian (normal) distribution (ND) and *t*-distribution (tD) of measured values x

The parameters of the normal distribution are:

- Mean value μ

$$\mu = \frac{\sum_{i=1}^n x_i}{n}. \quad (2.2.1-2)$$

- Variance σ^2

$$\sigma^2 = \frac{\sum_{i=1}^n (x_i - \mu)^2}{n}. \quad (2.2.1-3)$$

To avoid the scale effect, standardized values with

$$z = \frac{x - \mu}{\sigma} \quad (2.2.1-4)$$

are often used. Equation (2.2.1-1) is transformed into (2.2.1-5):

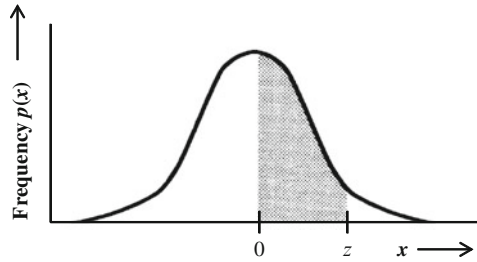
$$p(x) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{z^2}{2}\right). \quad (2.2.1-5)$$

Equation (2.2.1-5) holds true for the *standardized normal distribution*.

In the literature one can find some tables for z -values [7]. Table A.1 gives the areas between the boundary $z = 0$ and a chosen value z (see Fig. 2.2.1-3).

Because of the symmetry of the normal distribution the table gives p -values only for positive values of z . With this table one can ask, for example, what percentage of determinations will fall between two chosen boundaries (see Challenge 2.2.1-2).

Fig. 2.2.1-3 The shaded area describes the probability $p(x)$ of finding a value between 0 and z



In analytical practice, random samples of the basic population are investigated. The parameters μ and σ are substituted by the estimated values $\bar{x}$ and s for n measurements, which are calculated by (2.2.1-6) and (2.2.1-7), respectively:

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n}, \quad (2.2.1-6)$$

$$s = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n - 1}}. \quad (2.2.1-7)$$

Both values can be obtained by MS Excel functions = AVERAGE(Data) and =STDEV(Data), respectively.

Note the calculation of the mean value $\bar{x}$ as well as the standard deviation s is based on the normal distribution of the data set.

However, there are data sets for which no assumptions about the distribution of the population can be made. These data sets are handled by so-called *robust methods* [9, 10]. The central tendency is expressed by the *median* $\tilde{x}$ instead of the mean value $\bar{x}$. The median is resistant to outlying observations which have a large effect on the mean and the standard deviation.

After ranking the n data, the median $\tilde{x}$ is the middle value of the given numbers in ascending order.

The median of a ordered data set $x_1, x_2, \dots, x_n$ is

$$\tilde{x} = x_{\frac{n+1}{2}} \quad (2.2.1-8)$$

when the size of the distribution is *odd*, and

$$\tilde{x} = \frac{1}{2} \left(x_{\frac{n}{2}} + x_{\frac{n}{2}+1} \right) \quad (2.2.1-9)$$

when the size of the distribution is *even*.

In practice, the median is calculated by the Excel function = MEDIAN(Data) without ranking of the data set.

Challenge 2.2.1-1

The mean values of 40 batches of an intermediate product of a synthesis for an active pharmaceutical ingredient (API), calculated as the content relative to a standard, are given in Table 2.2.1-1.

Create the histogram for the data set with an appropriate number of classes!

Can the data set be considered normally distributed?

Table 2.2.1-1 Mean values $\bar{x}$ in % (w/w) of 40 batches of an intermediate product of a synthesis

n	$\bar{x}$	n	$\bar{x}$	n	$\bar{x}$	n	$\bar{x}$
1	103.9	11	96.2	21	99.7	31	96.9
2	102.7	12	99.9	22	100.6	32	108.0
3	101.0	13	92.3	23	107.5	33	105.8
4	94.8	14	101.2	24	90.5	34	94.6
5	105.2	15	100.8	25	108.8	35	102.8
6	100.4	16	99.0	26	101.9	36	104.2
7	97.0	17	100.8	27	102.5	37	99.9
8	101.6	18	104.0	28	97.4	38	106.4
9	109.0	19	99.2	29	107.0	39	103.5
10	90.8	20	109.7	30	104.5	40	96.7

Solution to Challenge 2.2.1-1

The mean values $\bar{x}$ arranged in increasing size are listed in Table 2.2.1-2.

With the rule of thumb for the choice of the number of classes $k = \sqrt{40} = 6.3$, seven classes are chosen. The number of mean values n_j which belong to the k classes are given in Table 2.2.1-3. The histogram is visualized in Fig. 2.2.1-4 from the data of Table 2.2.1-3. Figure 2.2.1-4 shows that the mean values $\bar{x}$ may be regarded as normally distributed, which is demonstrated by the bell-shaped curve in Fig. 2.2.1-5. (A statistical test for normal distribution is presented in Sect. 3.2.1.)

Table 2.2.1-2 Mean values $\bar{x}$ in % (w/w) arranged in increasing size

n	$\bar{x}$	n	$\bar{x}$	n	$\bar{x}$	n	$\bar{x}$
24	90.5	16	99.0	14	101.2	30	104.5
10	90.8	19	99.2	8	101.6	5	105.2
13	92.3	21	99.7	26	101.9	33	105.8
34	94.6	12	99.9	27	102.5	38	106.4
4	94.8	37	99.9	2	102.7	29	107.0
11	96.2	6	100.4	35	102.8	23	107.5
40	96.7	22	100.6	39	103.5	32	108.0
31	96.9	15	100.8	1	103.9	25	108.8
7	97.0	17	100.8	18	104.0	9	109.0
28	97.4	3	101.0	36	104.2	20	109.7

Table 2.2.1-3 Classes k with their width as well as the number of the mean values n_j for each class k

k	Width	n_j
1	90–93	3
2	93–96	2
3	96–99	6
4	99–102	12
5	102–105	8
6	105–108	6
7	108–111	3

Fig. 2.2.1-4 Histogram generated with the data of Table 2.2.1-3

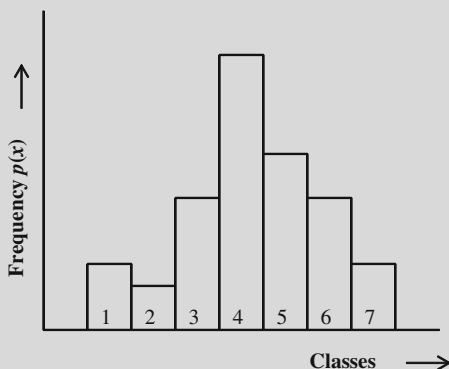
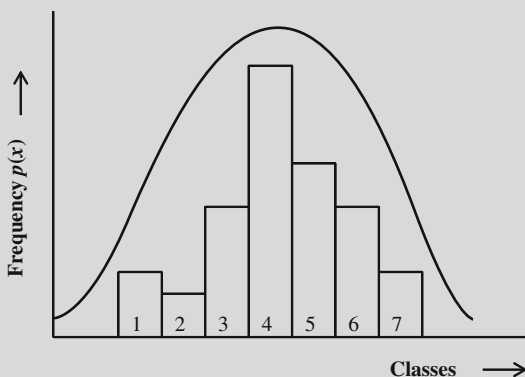


Fig. 2.2.1-5 Bell-shaped curve of the histogram in Fig. 2.2.1-4



Challenge 2.2.1-2

Calibration standards were prepared in the range 90–110% (w/w) for the determination of the content of the API with the same method as given in Challenge 2.2.1-1.

- What percentage of determinations will fall in this range?
- What percentage of determinations would fall in the range 99–101% (w/w)?

Solution to Challenge 2.2.1-2

According to the results in Challenge 2.2.1-1, the data set of the mean values $\bar{x}_i$ is normally distributed. Using the data set given in Table 2.2.1-2, the grand mean is $\bar{\bar{x}} = 101.2\%$ (w/w) and the standard deviation is $s = 4.857\%$ (w/w). These values are used for the calculation of the z -values according to (2.2.1-4).

(a) The z -value of the lower limit is $z_{ll} = -2.31$ and that of the upper limit is $z_{ul} = 1.81$. According to Table A.1 the probability p of finding an value between 0 and z is 0.4896 for the lower and 0.4649 for the upper limit, which gives the sum 0.9545. Thus, 95.5% of the results will fall within the range of the calibration standards, and only 4.5% will fall outside.

(b) For the range 99–101%(w/w), only 19.3% of the values are included and 80.7% fall outside, which is calculated by the intermediate quantities: $z_{ll} = -0.46$, $p(z_{ll}) = 0.1772$, $z_{ul} = -0.04$, $p(z_{ul}) = 0.0160$, $p(z_{ll} + z_{ul}) = 0.1932$.

Challenge 2.2.1-3

The screening of atrazine on a field by ELISA has yielded the mean values of 12 samples (n) obtained by triplicates given in Table 2.2.1-4.

Table 2.2.1-4 Data set obtained by screening of atrazine using ELISA

n	$\bar{x}_{\text{atrazine}}$ ppb (w/w)	n	$\bar{x}_{\text{atrazine}}$ ppb (w/w)	n	$\bar{x}_{\text{atrazine}}$ ppb (w/w)
1	2.5	5	4.6	9	13.8
2	0.9	6	0.5	10	1.2
3	1.1	7	8.6	11	0.8
4	7.9	8	3.1	12	6.4

(a) Calculate the mean value $\bar{x}$, the standard deviation s , and the median $\tilde{x}$; (b) Calculate the same parameters after addition of a further value 100 ppb (w/w) to the data set. Evaluate the results.

Solution to Challenge 2.2.1-3

(a) The mean value $\bar{x}$ and the standard deviation s calculated by (2.2.1-6) and (2.2.1-7), respectively, are $\bar{x} = 4.28$ ppm (w/w) and $s = 4.145$ ppm (w/w). In order to calculate the median, the data set has to be ordered (Table 2.2.1-5). Note that the median can also be calculated by the Excel function = MEDIAN(Data) without ordering of the data set.

Because the rank n is even the median is obtained by (2.2.1-9); the median is the mean of the values of rank 6 and 7: $\tilde{x} = 2.80$ ppm (w/w).

(continued)

Table 2.2.1-5 Analytical values of Table 2.2.1-4 in their ascending order n

n	x_{atrazine} ppb (w/w)	n	x_{atrazine} ppb (w/w)	n	x_{atrazine} ppb (w/w)
1	0.5	5	1.2	9	6.4
2	0.8	6	2.5	10	7.9
3	0.9	7	3.1	11	8.6
4	1.1	8	4.6	12	13.8

- (b) After addition of the value 100 ppb (w/w) to the data set the mean value is $\bar{x} = 11.65$ ppb (w/w) and the standard deviation is $s = 26.842$ ppb (w/w). Because the rank is now odd ($n = 13$) the median is the observation with rank $(13 + 1)/2 = 7$, according to (2.2.1-8): $\tilde{x} = 3.1$ ppb (w/w).

Whereas the addition of a single but outlying observation causes a large effect on the mean value as well as on the standard deviation, the median is hardly changed. The mean value increases from 4.28 to 11.65 ppb (w/w) whereas the median increases only from 2.8 to 3.1 ppb (w/w), which shows that the median is a better representative of the central tendency after addition of only one value to the data set.

2.2.2 Standard Deviation

The standard deviation s is calculated from n replicate measurements of the same sample by (2.2.1-7). The number of *degrees of freedom* is $df = n - 1$ which corresponds to the number of control measurements.

The standard deviation obtained from replicates of different samples with varying content is calculated by (2.2.2-1)

$$s = \sqrt{\frac{\sum_{i=1}^{n_A} \sum_{j=1}^m (x_{ij} - \bar{x}_i)^2}{n - m}} \quad (2.2.2-1)$$

with $df = n - m$, in which m is the number of samples, n_A is the number of replicates for each sample, and n is the total number of determinations, $n = m \cdot n_A$.

Equation (2.2.2-2) should be used for the computation of s :

$$s = \sqrt{\frac{\sum_{j=1}^m SS_i}{n - m}}. \quad (2.2.2-2)$$

SS_i is the sum of squares of the sample i , which is a calculator function and also a MS Excel function = DEVSQ(Data).

In the special case of paired replicates, each determination is carried out in duplicate. The standard deviation is calculated according to (2.2.2-3):

$$s = \sqrt{\frac{\sum (x'_j - x''_j)^2}{2 \cdot m}}. \quad (2.2.2-3)$$

The degrees of freedom $df = m$, x'_j and x''_j are the paired values of double measurements for each sample, and m is the number of samples.

The *variance* (*var*) is the square of the standard deviation:

$$\text{var} = s^2. \quad (2.2.2-4)$$

The *relative standard deviation* s_r is given by

$$s_r = \frac{s}{\bar{x}} \quad (2.2.2-5a)$$

and when it is expressed as a percentage by

$$s_r\% = 100 \cdot s_r \quad (2.2.2-5b)$$

which is, for example, an appropriate parameter for the comparison of precision of various analytical methods.

The *standard deviation of the means* (SEM) $s(\bar{x})$ is called the standard error of the mean and is calculated using the equation

$$s(\bar{x}) = \frac{s}{\sqrt{n}}. \quad (2.2.2-6)$$

The standard error of the mean is the standard deviation of the sample mean estimate of a population. It represents the variation associated with a mean value. The SEM is the expected value of the standard deviation of means of several samples.

Challenge 2.2.2-1

The process standard deviations for the determination of sulphur in steels according to the volumetric titration of SO_2 after burning of the samples were obtained by two different methods:

Method A: Repeated measurements of the *same* steel standard. The results are listed in Table 2.2.2-1.

(continued)

Table 2.2.2-1 Analytical values for a steel standard

Replicate	x in % (w/w)
1	0.0259
2	0.0238
3	0.0257
4	0.0242
5	0.0267
6	0.0239
7	0.0248
8	0.0259
9	0.0262
10	0.0241
11	0.0240

Table 2.2.2-2 Analytical values for ten steel standards

Standard	x' in % (w/w)	x'' in % (w/w)
1	0.0252	0.0236
2	0.0096	0.0110
3	0.0298	0.0282
4	0.0430	0.0448
5	0.0274	0.0281
6	0.0326	0.0294
7	0.0456	0.0480
8	0.0156	0.0135
9	0.0352	0.0330
10	0.0362	0.0374

Method B: Double measurements with ten *different* steel standards always of different content. The results are given in Table 2.2.2-2.

Calculate the standard deviation and give the degrees of freedom for both methods.

Note that the statistical test for normal distribution, the requirement for standard deviation, is given in Sect. 3.2.1.

Solution to Challenge 2.2.2-1

Method A: The standard deviations for the data set in Table 2.2.2-1 are calculated by (2.2.1-5), but this is a function on every hand calculator and an Excel function = STDEV(Data).

The standard deviation is $s = 0.0010\%$ (w/w) S, obtained by $df = 10$ degrees of freedom.

Method B: The standard deviation is $s = 0.00137\%$ (w/w) S calculated by (2.2.2-3) with the intermediate quantities $\sum (x' - x'')^2 = 0.0000375$ and $df = m = 10$.

Note that the degrees of freedom df are equal for both methods!

Challenge 2.2.2-2

An analytical laboratory has to determine manganese in steels with contents between 0.35 and 1.15% (w/w) Mn. For the determination of the standard deviation of the analytical method, five steel standards were analyzed by the volumetric method. The results are presented in Table 2.2.2-3.

Table 2.2.2-3 Analytical values for the five steel standards in % (w/w) Mn

Standard 1	0.31	0.30	0.29	0.32
Standard 2	0.59	0.57	0.58	0.57
Standard 3	0.71	0.69	0.71	0.71
Standard 4	0.92	0.92	0.95	0.95
Standard 5	1.18	1.17	1.21	1.19

Calculate the standard deviation.

Solution to Challenge 2.2.2-2

The standard deviation for the data set in Table 2.2.2-3 is calculated by (2.2.2-1). The sums of squares SS_i obtained by the MS Excel function = DEVSQ(Data) are:

Standard	1	2	3	4	5
SS_i	0.0005	0.000275	0.0003	0.0009	0.000875

The standard deviation is $s = 0.014\%$ (w/w) Mn which is obtained with $\sum SS_i = 0.00285$, $n = 20$, $m = 5$, and $df = 15$.

Note that the calculation of the standard deviation by (2.2.2-1) is only allowed if the variances of groups are homogeneous, which will be tested later (see Challenge 3.4-1).

2.2.3 Confidence Interval

Measured values which follow a normal distribution can occur in the whole range defined as $-\infty < x < \infty$. Therefore, it is useful to define dispersion ranges which include a certain number of measured values with a given high level of significance P , usually $P = 95\%$ or $P = 99\%$. The integration interval for $P = 95\%$ is $\mu \pm 1.96 \cdot \sigma$ and its limits are called *confidence limits* at the significance level $P = 95\%$. The range between the limits is called the *confidence interval*. Note that the integration between the limits $\pm 1.96 \cdot \sigma$ covers 95% of the values x_i . Thus, there is a probability of 95% that a measured value $\bar{x}$ will fall in the range $\mu \pm 1.96 \cdot \sigma$ under the assumption the values x_i of the mean belong to the same population.

For small sample sizes with n samples, the normal distribution $n(\sigma, \mu)$ is substituted by the t -distribution $n_t(s, \bar{x}, n)$. Figure 2.2.1-2 shows the relation between the normal distribution of a given population and the t -distribution for small samples. One can recognize that the t -distribution (curve tD) is broader at the base and the confidence interval is also broader. The confidence limits are given for the various n and degrees of freedom df, respectively, in the t -table (Table A.2) or by Excel function =TINV(α , df). Note that α is the risk, which is connected with the significance level by the relation $\alpha = 1 - P$.

Note that only two-tailed values are directly available from this function. In order to obtain a one-tailed critical value for the significance level α and df degrees of freedom the function =TINV(2α , df) is used. (One- and two-tailed values are explained in detail in Chap. 3.)

The confidence interval is calculated by (2.2.3-1):

$$\Delta\bar{x} = \frac{s_x \cdot t(P, \text{df})}{\sqrt{n}}. \quad (2.2.3-1)$$

The t -values, i.e. the critical values of the t -distribution, are taken from Student's t -table for a certain significance level P (usually 95 or 99%) and the degrees of freedom df refers to the data set from which the standard deviation s_x is obtained.

The analytical result is expressed in the form:

$$\bar{x} \pm \Delta\bar{x}, \text{ given in the units of measurement.} \quad (2.2.3-2)$$

The mean value $\bar{x}$ is calculated by (2.2.1-2). But there is still a question: how big may the difference of two or more measurement values be for the formation of the mean value? Can all values x_i obtained be utilized or there are limits?

The *critical difference* Δ_{crit} between the highest and the lowest measurement values in a set of repeated determinations is given by *Pearson's criterion*:

$$\Delta_{\text{crit}} = |x_{\text{max}} - x_{\text{min}}| < D(P, n_j) \cdot s_x. \quad (2.2.3-3)$$

The Pearson factors $D(P, n_j)$ for $P = 95\%$ and the number of repeated determinations n_j are given in Table 2.2.3-1.

Table 2.2.3-1 Pearson factors $D(P, n_j)$ for the critical difference between the highest and lowest measurement value of repeated determinations n_j with the significance level $P = 95\%$

n_j	2	3	4
$D(P, n_j)$	2.77	3.31	3.65

For example, the difference between the two measurement values may not exceed the limit $2.77 \cdot s_x$ for a double determination.

However, sometimes the simple relation $\Delta_{\text{crit}} < 2 \cdot s_x$ is used; therefore, the limit criterion used should be given in the documents.

Challenge 2.2.3-1

Let us come back to the determination of sulphur in steel (Challenge 2.2.2-1).

Calculate the confidence interval for the mean value of sulphur using method A and method B at the significance level $P = 95\%$ for

(a) Double determinations

(b) Fourfold determinations

Solution to Challenge 2.2.3-1

The confidence interval is calculated by (2.2.3-1).

The results are listed in Table 2.2.3-2. The values of s_x and df were determined in Challenge 2.2.2-1.

Table 2.2.3-2 Intermediate quantities and results of the calculation of the confidence interval $\Delta\bar{x}$ for the determination of sulphur in steel by different methods

Parameter	Method A	Method B
s_x	0.0011 in % (w/w)	0.0014 in % (w/w)
df	10	10
$t(P = 95\%, \text{df})$	2.228	2.228
a. $\Delta\bar{x}$ for $n = 2$	0.0017 in % (w/w)	0.0022 in % (w/w)
b. $\Delta\bar{x}$ for $n = 4$	0.0012 in % (w/w)	0.0015 in % (w/w)

Challenge 2.2.3-2

(a) According to the procedure given in Challenge 2.2.2-2, the double determination of manganese in a steel sample yields the values $x_1 = 0.65\%$ (w/w) Mn and $x_2 = 0.63\%$ (w/w) Mn.

Present the analytical result in the form $x \pm \Delta\bar{x}\%$ (w/w) Mn.

Give a verbal interpretation of the result.

(b) The analytical results obtained with triplicates of a sample of manganese steel are:

% (w/w) Mn	0.65	0.63	0.68
------------	------	------	------

Test whether the calculation of the mean value is permitted.

What should one do if the limit is exceeded?

Solution to Challenge 2.2.3-2

- (a) The confidence interval is $\Delta\bar{x} = 0.021\%$ (w/w) Mn calculated for $n_j = 2$ (double determination) with the data obtained by Challenge 2.2.2-2: $s_x = 0.014\%$ (w/w) Mn and $t(P = 95\%, df = 15) = 2.131$.

The analytical result is $0.64 \pm 0.02\%$ (w/w) Mn.

The true value of the content of manganese in the steel sample lies in the range 0.62–0.66% (w/w) Mn. But this is true only for the significance level $P = 95\%$, with the risk $\alpha = 5\%$ that the true value will lie outside this range.

- (b) For $n_j = 3$ the Pearson factor is 3.31. With $s_x = 0.014\%$ (w/w) Mn, the critical difference is 0.046% (w/w) Mn, but the difference in the experimental values is $x_{\max} - x_{\min} = 0.05\%$ (w/w) Mn. The calculation of the mean value is not permitted.

One should at best make a further analysis.

2.2.4 Confidence Interval and Quality

The quality control of products in environmental compartments and elsewhere requires decisions on the basis of analytical results, which means deciding whether a limit value is transgressed or not. Such a limit or threshold value stipulated in official documents can be an upper limit (e.g. in the case of environmental compartments) or a lower limit (e.g. for the potassium content of a fertilizer).

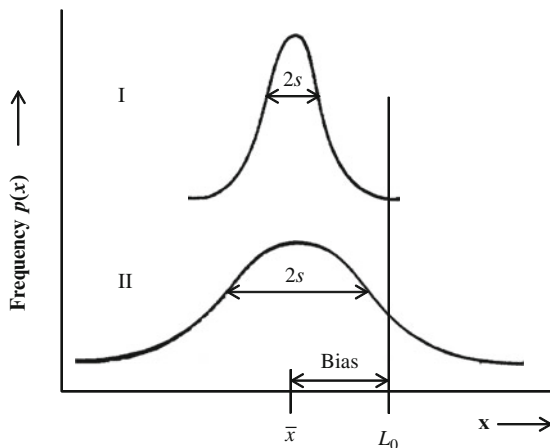
Let us take an example: the specified threshold for the content of the monomer styrene in industrially produced polystyrene for a certain application is $\leq 0.8\%$ (w/w). Analytical quality assurance yields a content of 0.75% (w/w) for a batch of polystyrene. Is the limit value exceeded or not, i.e. is this batch has to be discarded or is the quality standard fulfilled? How is it to be recognized easily: this decision is of great economic interest?

But, as Fig. 2.2.4-1 shows, the decision cannot be made without knowledge of the confidence interval of the analytical result.

The same mean value $\bar{x}$ was obtained with two methods which are different in regard to their precision. The quality criterion is fulfilled in the upper case I, because the limit value L_0 falls outside the confidence interval CI. One says that L_0 does not belong to the parent basic population of the sample. But in case II with the larger standard error, L_0 is included in the basic population of the sample which means, in a statistical sense, there is no difference between $\bar{x}$ and L_0 . Therefore, the limit value is exceeded and, for example, the product cannot be delivered for sale.

For the control of limiting values, as well as some other problems, only the *one-sided* limit of the confidence interval is important. This is the upper limit in the

Fig. 2.2.4-1 Influence of the precision s on the transgression of the threshold value L_0 for the same analytical result $\bar{x}$



case of Fig. 2.2.4-1. The significance level of one-sided confidence intervals is also taken from the t -table or the MS Excel spreadsheet, but with another value for the statistical significance level. It is worth knowing that for the usual significance levels $t(\bar{P}_{\text{one-sided}} = 95\%, \text{df}) \approx t(P_{\text{two-sided}} = 90\%, \text{df})$.

An analytical mean value fulfils the quality standard for a required maximal threshold value L_0 if

$$\bar{x} + \frac{s \cdot t(\bar{P}_{\text{one-sided}}, \text{df})}{\sqrt{n_a}} \leq L_0. \quad (2.2.4-1)$$

The degrees of freedom df refer to the number of replicates with which the standard deviation of the analytical method s has been determined, and n_a is the number of replicates in the routine analysis.

As Figure 2.2.4-1 reveals, an analytical method with a small confidence interval is desirable because the experimentally determined mean value can be closer to the limit value without it being exceeded. If one inspects (2.2.3-1), the confidence interval for a given significance level, usually $P = 95\%$, is determined by the standard deviation of the method s and the degrees of freedom df for its determination as well as the number n_a of the replicates in the routine analysis. The larger the number n the smaller will be the value $\Delta\bar{x}$. But the influence of n on the value of the confidence interval falls exponentially, as demonstrated in Fig. 2.2.4-2 [9]. Many replicates in routine quality control quickly increases costs, but the effect is only small. Double determinations are often sufficient.

However, the standard deviation of the analytical method s is direct proportional to $\Delta\bar{x}$. Thus, it has the biggest influence on the magnitude of $\Delta\bar{x}$. The determination of s is a unique procedure, and therefore a larger number of replicates should be made. On the other hand, the greater the number of replicates for the determination of s , the smaller the t -value.

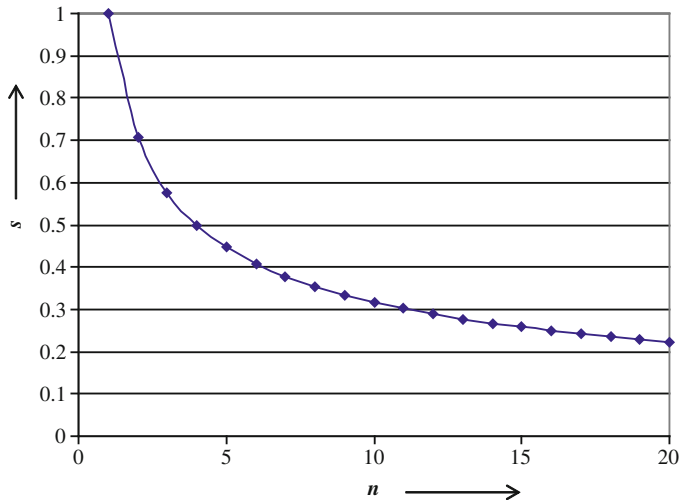


Fig. 2.2.4-2 Relation between s and the number of repeated measurements n [9]

Challenge 2.2.4-1

According to a company specification the content of benzene (bz) in technical n -hexane may not be greater than $L_0 = 0.80\%$ (v/v). The analytical quality control will be carried out by GC with n -octane as an internal standard (IS).

The process standard deviation was determined with varying numbers of replicates of the same sample:

Method A: 12 individual samples

Method B: 6 individual samples.

The relative peak areas $A_{\text{bz}}/A_{\text{IS}}$ obtained from the chromatograms are given in Table 2.2.4-1.

(continued)

Table 2.2.4-1 The relative peak areas $A_{\text{bz}}/A_{\text{IS}}$ obtained from the chromatograms

Replicate	$A_{\text{bz}}/A_{\text{IS}}$	Replicate	$A_{\text{bz}}/A_{\text{IS}}$
Method A			
1	0.855	7	0.866
2	0.834	8	0.873
3	0.862	9	0.819
4	0.860	10	0.854
5	0.854	11	0.886
6	0.843	12	0.875
Method B			
1	0.788	4	0.796
2	0.772	5	0.747
3	0.769	6	0.758

A_{bz} is the peak area of benzene and A_{IS} is the peak area of the internal standard n -octane

Which mean value of benzene $\bar{x}_{\text{bz}}$ may not be exceeded if

- (a) Double determinations or
- (b) Fourfold determinations will be carried out in the quality control?

Evaluate the results.

Solution to Challenge 2.2.4-1

The critical mean value of benzene $\bar{x}_{\text{crit,bz}}$ which may not be exceeded is calculated according to (2.2.4-1):

$$\bar{x}_{\text{crit,bz}} \leq L_0 - \frac{s \cdot t(\bar{P}_{\text{one-sided}}, \text{df})}{\sqrt{n_a}}$$

(2.2.4-2)

The intermediate quantities and the critical mean value of benzene $\bar{x}_{\text{crit,bz}}$ calculated according to (2.2.4-2) are given in Table 2.2.4-2 for the various conditions.

Table 2.2.4-2 Intermediate quantities and the limit value of benzene $\bar{x}_{\text{crit,bz}}$ calculated according to (2.2.4-2)

Parameter	Method A	Method B
Determination of the standard deviation s		
n	12	6
df	11	5
s in % (v/v)	0.0184	0.0182
$t(\bar{P}_{\text{one-sided}} = 95\%, \text{df})$	1.796	2.015
Routine quality control		
n_a	2	2
$\Delta \bar{x}$ in % (v/v)	0.023	0.026
$\bar{x}_{\text{crit,bz}}$ in % (v/v)	0.777	0.774
n_a	4	4
$\Delta \bar{x}$ in % (v/v)	0.017	0.018
$\bar{x}_{\text{crit,bz}}$ in % (v/v)	0.783	0.782

As Table 2.2.4-2 shows, the critical mean value of benzene $\bar{x}_{\text{crit,bz}}$ differs only minimally with the various conditions. Double determinations in the routine quality control and determination of the standard deviation of the analytical method with twelve replicates yields the critical value $\bar{x}_{\text{crit,bz}} = 0.78\%$ (v/v). Increasing the numbers of replicates for the determination of s as well as in the routine quality control does not have a practical influence on the critical mean value.

Challenge 2.2.4-2

A company produces polystyrene for a certain application. The content of the residual monomer may not exceed 0.60% (w/w) styrene. The monomer will
(continued)

Table 2.2.4-3 Analytical results x in % (w/w) styrene obtained by two replicates with six polystyrene samples by MHE-HS-GC

Sample	1	2	3	4	5	6
x'	0.573	0.654	0.916	0.439	0.753	0.848
x''	0.525	0.691	0.972	0.489	0.812	0.892

be analyzed by MHE-HS-GC (see Chap. 7). The standard deviation of the analytical method was determined by two replicates with six samples. The results are listed in Table 2.2.4-3.

In the routine quality control of a sample the following analytical results were obtained:

% (w/w) styrene	0.562	0.591	0.559
-----------------	-------	-------	-------

Will the sample meet the quality requirement?

Solution to Challenge 2.2.4-2

The standard deviation of the analytical method s_x calculated according to (2.2.2-3) with $\sum (x' - x'')^2 = 0.014726$ $(\%(\text{w/w}))^2$ and $m = 6$ is $s = 0.03503\%$ (w/w).

The confidence interval calculated by (2.2.3-1) with $\bar{x} = 0.5707\%$ (w/w), $t(\bar{P}_{\text{one-sided}} = 95\%, \text{df} = 6) = 1.943$, and $n_a = 3$ is $\Delta\bar{x} = 0.0393\%$ (w/w). Thus, the upper confidence limit is $\bar{x} + \Delta\bar{x}_{\text{one-sided}} = 0.61\%$ (w/w) styrene. This value exceeds the documented quality limit of $L_0 = 0.60\%$ (w/w), and therefore the sample does not fulfil the quality requirements. It cannot be delivered for sale.

2.2.5 Propagation of Errors

When the final result is obtained from more than one independent measurement, or when it is influenced by two or more independent sources of errors, these errors can be accumulated or compensated. This is called the propagation of errors.

In the case of independent variables $x_1, x_2, \dots, x_n$, i.e. if there is no correlation between the x -values, i.e. the covariances

$$\text{cov}(x_1, x_2, \dots, x_n) = \frac{1}{n-1} \cdot \left[\sum (x_{1i} - \bar{x}_1) \cdot (x_{2i} - \bar{x}_2) \cdot \dots \cdot (x_{ni} - \bar{x}_n) \right] = 0, \quad (2.2.5-1)$$

the total error can be estimated according to the Gaussian law of error propagation:

$$\sigma_x^2 = \left(\frac{\partial f}{\partial x_1}\right)^2 \cdot \sigma_{x_1}^2 + \left(\frac{\partial f}{\partial x_2}\right)^2 \cdot \sigma_{x_2}^2 + \cdots + \left(\frac{\partial f}{\partial x_n}\right)^2 \cdot \sigma_{x_n}^2. \quad (2.2.5-2)$$

For addition or subtraction the variances are additive:

$$\sigma_x^2 = \sigma_{x_1}^2 + \sigma_{x_2}^2 + \cdots + \sigma_{x_n}^2. \quad (2.2.5-3)$$

For multiplication or division the squared relative standard deviations are additive:

$$\left(\frac{\sigma_x}{x}\right)^2 = \left(\frac{\sigma_{x_1}}{x_1}\right)^2 + \left(\frac{\sigma_{x_2}}{x_2}\right)^2 + \cdots + \left(\frac{\sigma_{x_n}}{x_n}\right)^2. \quad (2.2.5-4)$$

Note that, as mentioned above, these equations are correct only when the variables are independent, i.e. if they are not correlated!

Challenge 2.2.5-1

The content of a pharmaceutical product will be detected by HPLC.

The percentage content of the active pharmaceutical ingredient (API) $x_{\text{API}}\%$ (w/w) is calculated by (2.2.5-5).

$$x_{\text{API}}\% \text{ (w/w)} = \frac{\bar{A}_s \text{ in counts} \cdot 100}{(c_s \text{ in g L}^{-1})(\text{Rf in counts L g}^{-1})} \quad (2.2.5-5)$$

$\bar{A}_s$ is the mean peak area of the sample obtained by the chromatogram, c_s is the concentration of the sample, and Rf is the response factor which is determined with a solution of chemical reference substance (CRS) according to (2.2.5-6):

$$\text{Rf} = \frac{\bar{A}_{\text{CRS}} \text{ in counts}}{(c_{\text{CRS}} \text{ in g L}^{-1}) \cdot (x_{\text{CRS}} \text{ in \% (w/w)}) \cdot 0.01}. \quad (2.2.5-6)$$

$\bar{A}_{\text{CRS}}$ is the mean of the peak area of CRS, c_{CRS} is the concentration of CRS, and $x_{\text{CRS}}\%$ (w/w) is the certified content of CRS.

According to the United States Pharmacopeia (USP) the relative standard deviation of the precision of injection of the sample should be $s_r\% \leq 1.0$.

Testing the precision of injection, as usual in pharmaceutical analyses, a sample CRS was measured with six replicates. The peak areas obtained from the HPLC chromatograms are presented in Table 2.2.5-1.

The experimental data for the determination of the content of the API $x_{\text{API}}\%$ (w/w) are given in Table 2.2.5-2.

Table 2.2.5-1 Peak areas A obtained from the HPLC chromatograms of a CRS solution	Replicate	A in counts
	1	678,458
	2	670,554
	3	678,458
	4	664,119
	5	680,246
	6	672,179

Table 2.2.5-2 Experimental data for determination of the API	Solutions	
	Determination of Rf	$c_{\text{CRS}} = 0.813 \text{ g L}^{-1}$
	Determination of x_{API}	$c_{\text{s}} = 0.803 \text{ g L}^{-1}$
	Certified content of the CRS x_{CRS}	99.15% (w/w)
	Peak areas A in counts obtained by the HPLC chromatograms	
	Rf	Sample
	114,856	112,969
	115,681	111,781
	114,836	111,876
	113,592	113,006

(a) Test whether the claimed precision of injection is achieved;

(b) Calculate the API content of the sample with its confidence interval $\bar{x} \pm \Delta x\%$ (w/w) API.

Solution to Challenge 2.2.5-1

(a) The data set of Table 2.2.5-1 gives $s = 6,190.1$ counts, $\bar{A}_{\text{CRS}} = 674,002.3$ counts, and $s_{\text{r}}\% = 0.92$.
The relative standard deviation $s_{\text{r}}\%$ is smaller than the limit value given in USP. This means the injection precision of the HPLC method is achieved.

(b) The intermediate quantities are:

– Rf = $142,343.1 \text{ counts L g}^{-1}$ calculated by (2.2.5-6) with $\bar{A}_{\text{Rf}} = 114,741.3$ counts and further data given in Table 2.2.5-2

– $s_{\text{A}_{\text{Rf}}}^2 = 742,016.9 \text{ counts}^2$,

– $\bar{A}_{\text{s}} = 112,408.0$ counts,

– $s_{\text{A}_{\text{s}}}^2 = 449,492.7 \text{ counts}^2$.

The content of the sample calculated by (2.2.5-5) is $\bar{x}_{\text{API}}\%$ (w/w) = 98.34.
The total variance s_{tot}^2 for the determination of the API derived according to (2.2.5-2) is

(continued)

$$s_{\text{tot}}^2 = \left(\frac{100}{c_s \cdot \text{Rf}} \right)^2 s_{A_s}^2 + \left(\frac{100 \cdot \bar{A}_s}{c_s \cdot (\text{Rf})^2} \right)^2 s_{A_{\text{Rf}}}^2. \quad (2.2.5-7)$$

s_{tot}^2 calculated by (2.2.5-7) is $s_{\text{tot}}^2 = 0.3576 \text{ g}^2 \text{ L}^{-2}$ and the standard deviation is $s_{\text{tot}} = 0.5980 \text{ g L}^{-1}$, respectively.

The confidence interval

$$\Delta \bar{x}_{\text{API}} \% (\text{w/w}) = \frac{s_{\text{tot}} \cdot t(P, \text{df})}{\sqrt{n}} \quad (2.2.5-8)$$

is $\Delta \bar{x} \% (\text{w/w}) = 0.73$ calculated with $\text{df}_{\text{total}} = \text{df}_{\text{RF}} + \text{df}_s = 6$, $t(P = 95\%, \text{df} = 6) = 2.447$, and $n = 4$.

Result: The content of the sample is $\Delta \bar{x} \% (\text{w/w}) = 98.34 \pm 0.73$. The true value lies in the range 97.61–99.07% (w/w) at the significance level $P = 95\%$ and with the risk $\alpha = 5\%$ that the true value may be found outside this range.

Challenge 2.2.5-2

Let us now estimate the errors in photometric analysis which is an important method in AQA. As example, we will choose IR spectrophotometric analysis which must often be applied in AQA (see for example Challenge 3.3-3).

The spectrophotometric analysis is based on *Lambert–Beer's law*

$$A = \alpha \cdot c \cdot l \quad (2.2.5-9)$$

where α is the absorptivity (a constant which is usually given in $\text{L mol}^{-1} \text{ cm}^{-1}$ or in $\text{m}^2 \text{ mol}^{-1}$), c is the concentration in mol L^{-1} , and l is the optical path length, i.e. the diameter of the cuvette.

In IR spectrophotometry the optical path length lies in the μm range, and therefore it is determined by the interference method. The order of the interferences n which are obtained if the empty cuvette is traversed by IR light is calculated by

$$r = 2 \cdot l \cdot v_{\text{max}} \quad (2.2.5-10)$$

where v_{max} is the maximum of the interference, r is the number of the reflection, and l is the optical path length. According to (2.2.5-10) l is obtained from the slope of the function $n = f(2v)$.

Using standard calibration solutions with amount m in volume V , the absorptivity α is calculated by (2.2.5-11). Usually, in IR spectrophotometry the constant α is given in the units $\text{L g}^{-1} \text{ cm}^{-1}$ according to (2.2.5-11):

(continued)

$$\alpha = \frac{A \cdot V}{l \cdot m}, \quad (2.2.5-11)$$

which will also be used in the following.

Finally, let us turn to the absorbance A which is measured by the spectrophotometer.

The absorbance is defined by

$$A = \frac{\log I_0}{\log I} = \log I_0 - \log I, \quad (2.2.5-12)$$

where I_0 and I are the intensity of the reference beam and the intensity of the sample beam, respectively.

The error of the measurement of the absorbance is given by the law of error propagation according to (2.2.5-2):

$$\sigma_A^2 = \left(\frac{\partial A}{\partial I_0} \right)^2 \cdot \sigma_{I_0}^2 + \left(\frac{\partial A}{\partial I} \right)^2 \cdot \sigma_I^2. \quad (2.2.5-13)$$

From (2.2.5-12) follows

$$\sigma_A^2 = \left(\frac{\partial(\log I_0 - I)}{\partial I_0} \right)^2 \cdot \sigma_{I_0}^2 + \left(\frac{\partial(\log I_0 - I)}{\partial I} \right)^2 \cdot \sigma_I^2, \quad (2.2.5-14)$$

which gives (2.2.5-15), and with $\log e = 0.43$ (2.2.5-16), respectively

$$\sigma_A^2 = \left(\frac{\log e}{I_0} \right)^2 \cdot \sigma_{I_0}^2 + \left(\frac{\log e}{I} \right)^2 \cdot \sigma_I^2 \quad (2.2.5-15)$$

$$\sigma_A^2 = \left(\frac{0.43}{I_0} \right)^2 \cdot \sigma_{I_0}^2 + \left(\frac{0.43}{I} \right)^2 \cdot \sigma_I^2. \quad (2.2.5-16)$$

The Challenges are:

- Derive the equation for variance of the absorptivity σ_α^2 from the law of propagation of errors.
- A problem in spectrophotometry is the magnitude of the chosen *absorbance* A . Decide if the relative error of the measurement of the absorbance is constant or variable. Derive the relation for this relative error and create a graph for the relative error of the measurement of the absorbance using values in the range 0.025–2.5 in appropriate steps. Estimate the result with regard to the choice of an optimal range for the measurement of A .

(continued)

- (c) A further parameter in IR spectrophotometry is the magnitude of the slit width and the *slit width program*, respectively. Decide which slit width program should be chosen.
- A tip: consider the fact that I_0 grows with the square of the slit width.
- (d) Calculate the absorptivity α in $\text{L cm}^{-1}\text{g}^{-1}$ with its random error for the carbonyl band of lactic acid ester (LAE) at $1,735\text{ cm}^{-1}$ from the following data:

Sample solution	$m = 50.28\text{ mg LAE in } 10\text{ mL } n\text{-hexane}$		
Absorbance A	0.391525	0.391701	0.392668
	0.393124	0.392147	0.391010

The diameter of the cuvette is determined by the interference maxima given in Table 2.2.5-4.

The random errors of mass m , diameter of the cuvette l , and volume V are estimated from the data sets given in Tables 2.2.5-3 and 2.2.5-5.

Note that the influence of temperature, the tolerance of the volumetric flask, and other factors are neglected here. This is the subject of Chap. 10. Calculate the percentage of the individual variances in the variance of the absorptivity.

- (e) Test whether a fivefold increase in sample volume will appreciably diminish the random error σ_A . The procedure for the determination of
(continued)

Table 2.2.5-3 Estimation of the random error of the balance

Number	Gross weight in g	Tare weight in g	Number	Gross weight in g	Tare weight in g
1	6.19740	6.09748	6	6.13155	6.03159
2	6.09595	5.99596	7	6.22193	6.12196
3	6.13175	6.03178	8	6.09995	6.00000
4	6.13467	6.03472	9	6.07420	5.97420
5	6.06939	5.96935	10	6.08567	5.98577

Table 2.2.5-4 Estimation of the random error of the diameter of the cuvette l from interference maxima $I_{\max}$ measured with the empty cuvette

Order number r	$I_{\max}$	Order number r	$I_{\max}$
1	793	11	1,128
2	829	12	1,160
3	861	13	1,191
4	895	14	1,226
5	927	15	1,259
6	960	16	1,292
7	993	17	1,327
8	1,027	18	1,358
9	1,060	19	1,394
10	1,092	20	1,425

Table 2.2.5-5 Estimation of the random error of the volume $V = 10$ mL. A 10 mL volumetric flask was filled up with water and the mass was determined.

Number	m in g	Number	m in g
1	9.964761	6	9.962722
2	9.974138	7	9.989396
3	9.983647	8	9.972522
4	9.985056	9	9.983176
5	9.997446	10	9.969295

Table 2.2.5-6 Estimation of the random error of the volume with $V = 50$ mL. A 50 mL volumetric flask was filled up with water and the mass was determined.

Number	m in g	Number	m in g
1	50.063150	6	50.061528
2	50.090792	7	50.116459
3	50.051704	8	50.116849
4	50.066276	9	50.031270
5	50.052144	10	50.097729

volume error is the same as for 10 mL volumes (Table 2.2.5-5). The experimental values are listed in Table 2.2.5-6.

Solution to Challenge 2.2.5-2

- (a) Using the law of propagation of errors, equation (2.2.5-2), the error of the absorptivity α is given by (2.2.5-17):

$$\sigma_{\alpha}^2 = \underbrace{\left(\sigma_A \cdot \frac{V}{l \cdot m}\right)^2}_{\text{I}} + \underbrace{\left(\sigma_V \cdot \frac{A}{l \cdot m}\right)^2}_{\text{II}} + \underbrace{\left(\sigma_l \cdot \frac{-A \cdot V}{l^2 \cdot m}\right)^2}_{\text{III}} + \underbrace{\left(\sigma_m \cdot \frac{-A \cdot V}{l \cdot m^2}\right)^2}_{\text{IV}} \quad (2.2.5-17)$$

Term I represents the error in the measurement of the absorbance, term II that for the volume of the measuring solution, term III that of the optical path length, and term IV that for the weight of the mass of the sample for the preparation of the measuring solution.

- (b) The equation for the relative error in the measurement of the absorbance σ_A/A is obtained from (2.2.5-15) with $\sigma_{I_0} = \sigma_I = 1$:

$$\frac{\sigma_A}{A} = \frac{\lg e}{A} \cdot \sqrt{\frac{1}{I_0^2} + \frac{1}{I^2}} \quad (2.2.5-18)$$

Next, I is substituted as follows:

From (2.2.5-12) one obtains for I

$$\log I = \log I_0 - A \quad (2.2.5-19)$$

(continued)

and

$$I = 10^{(\log I_0) - A} = \frac{I_0}{10^A}. \quad (2.2.5-20)$$

Equation (2.2.5-20) in (2.2.5-18) gives:

$$\frac{\sigma_A}{A} = \frac{\log e}{A} \cdot \sqrt{\frac{1}{I_0^2} + \frac{10^{2A}}{I_0^2}} = \frac{\log e}{I_0} \cdot \frac{\sqrt{1 + 10^{2A}}}{A}. \quad (2.2.5-21)$$

Equation (2.2.5-21) shows that the relative error for the measurement of the absorbance is not constant, but is a function of the absorbance.

The relative errors for the measurement of the absorbance for a chosen data set calculated by (2.2.5-21) with $I_0 = 1$ and $\log e = 0.43$ are listed in Table 2.2.5-7 and the graph is presented in Fig. 2.2.5-1.

As the graph shows, the minimum of the relative errors of the measurement of the absorbance is in the range from about 0.3 to about 1.0. Therefore, this is an optimal range for spectrophotometry. For solutions with absorbance lower than 0.1 or greater than 2.0 the relative error rises rapidly.

- (c) As (2.2.5-18) shows, the relative error of the measurement of the absorbance diminishes with increasing I_0 . Because I_0 grows with the square of the slit width, one should take the largest slit width or the split program.
- (d) The diameter of the cuvette l is determined by the regression analysis of the data set in Table 2.2.5-4 and is the slope of the function $n = f(2v)$ according to (2.2.5-10), $l = 0.01505$ cm. The standard error of the diameter corresponds to the standard deviation of the slope calculated by (5.2-13), which is $\sigma_l = 1.833 \times 10^{-5}$ cm obtained with $SS_{xx} = 2,934,826.2$ and $s_{y,x} = 0.0314027$.

(continued)

Table 2.2.5-7 Relative errors of the measurement of the absorbance $s_{r,A} = \sigma_A/A$ calculated by (2.2.5-21)

A	$\frac{\sigma_A}{A}$	A	$\frac{\sigma_A}{A}$	A	$\frac{\sigma_A}{A}$
	A		A		A
0.020	31.13	0.400	2.91	0.850	3.62
0.025	25.06	0.450	2.86	0.900	3.83
0.050	12.93	0.500	2.85	0.950	4.06
0.100	6.91	0.550	2.88	1.000	4.32
0.150	4.96	0.600	2.94	1.250	6.13
0.200	4.03	0.650	3.03	1.500	9.07
0.250	3.51	0.700	3.14	1.750	13.82
0.300	3.20	0.750	3.27	2.000	21.50
0.350	3.01	0.800	3.43	2.250	33.99

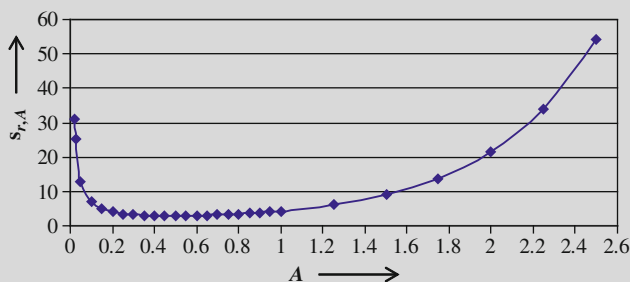


Fig. 2.2.5-1 Relative errors of the measurement of the absorbance $s_{r,A}$ as a function of the absorbance A

The mean value of the absorbance $\bar{A}$ is 0.392029. The absorptivity α is calculated according to (2.2.5-11):

$$\alpha = \frac{0.392029 \cdot 0.01 \text{ L}}{0.01505 \text{ cm} \cdot 0.05028 \text{ g}} = 5.18 \text{ L cm}^{-1} \text{ g}^{-1}. \quad (2.2.5-22)$$

The standard deviation of the measurement of the absorbance is $\sigma_A = 7.773 \times 10^{-4}$.

The standard deviation of the net values (difference between gross and tare of the data in Table 2.2.5-3) is $\sigma_m = 3 \cdot 98 \cdot 10^{-5} \text{ g}$.

The standard deviation of the filling of the volumetric flask is calculated using the data set in Table 2.2.5-5: $\sigma_V = 0.01128 \text{ mL}$ for a 10 mL volumetric flask.

The random error of the absorptivity α is calculated by (2.2.5-17) with the parameters $V = 0.01 \text{ L}$, $m = 0.05028 \text{ g}$, $l = 0.01505 \text{ cm}$, $A = 0.392029$, $\sigma_V = 1.128 \cdot 10^{-5} \text{ L}$, $\sigma_m = 3.979 \cdot 10^{-5} \text{ g}$, $\sigma_l = 1.833 \cdot 10^{-5} \text{ cm}$, and $\sigma_A = 0.0007773$.

The result is $\sigma_\alpha^2 = 0.0001962$ and $\sigma_\alpha = 0.0140$.

The absorptivity α calculated according to (2.2.5-11) is $\alpha = 5.18 \text{ L g}^{-1} \text{ cm}^{-1}$. Thus, the relative standard deviation calculated by (2.2.2-5a) is $s_r\% = 0.27$.

The percentages of the individual variances in the total variance σ_α^2 are given in Table 2.2.5-8.

As Table 2.2.5-8 shows, the measurement of the absorbance has the greatest influence on the random error of the absorptivity α , but one has to consider that all the uncertainties of Type B were rejected here.

- (e) The standard deviation of the filling of the volumetric flask is calculated with the data set in Table 2.2.5-6: $\sigma_V = 2.908 \cdot 10^{-5} \text{ L}$ for a 50 mL volumetric flask. The variance of the absorptivity is $\sigma_\alpha^2 = 0.000155$ calculated with the fivefold-increased mass $m = 2.514 \text{ g}$ according to
(continued)

Table 2.2.5-8 Percentages of the individual variances in the total variance of the absorptivity σ_x^2

σ_m^2 mass	σ_V^2 volume	σ_l^2 cuvette	σ_A^2 absorbance
8.6%	17.4%	20.3%	53.8%

(2.2.5-17). The standard deviation of the absorptivity is $\sigma_x = 0.01245$ and $s_r\% = 0.24$.

The increase in sample volume does not improve the precision in practice, but would only incur higher costs for the sample and the solvent.

References

1. ISO 3534-2 (2006) Statistics – vocabulary and symbols – part 2: applied statistics. International Organization for Standardization, Geneva
2. Joint Committee for Guides in Metrology (JCGM) (2008) International vocabulary of metrology – basic and general concepts and associated terms (VIM). International Organization for Standardization, Geneva
3. ISO Guide 98 (1995) Guide to the expression of uncertainty in measurement. International Organization for Standardization, Geneva
4. Danzer K (2007) Analytical chemistry – theoretical and metrological fundamentals. Springer, Berlin
5. Menditto A, Patria M, Magnusson B (2007) Understanding the meaning of accuracy, trueness and precision. *Accred Qual Assur* 12:45–47
6. Danzer K (1989) Robuste Statistik in der analytischen Chemie. *Fresenius Z Anal Chem* 335:869–875
7. Massart DL, Vandeginste BGM, Buydens LMC, De Jong S, Lewi PJ, Smeyers-Verbeke J (1997) Handbook of chemometrics and qualimetrics – part A. Elsevier, Amsterdam
8. Ellison SLR, Berwick VJ, Duguid Farrant TJ (2009) Practical statistics for the analytical scientist, 2nd edn. RSC, Cambridge
9. Doerffel K (1990) Statistik in der analytischen Chemie, 5 Aufl. Deutscher Verlag für Grundstoffindustrie, Leipzig
10. Danzer K (1989) Robuste Statistik in der analytischen Chemie. *Fresenius Z Anal Chem* 335:869–875

Challenges in Analytical Quality Assurance

Reichenbächer, M.; Einax, J.W.

2011, XIX, 356 p., Hardcover

ISBN: 978-3-642-16594-8