



OpenLab CDS

Data Analysis Reference Guide

Notices

Document Information

Document No: D0013815 Rev. A
Edition: 04/2022

Copyright

© Agilent Technologies, Inc. 2012-2022

No part of this manual may be reproduced in any form or by any means (including electronic storage and retrieval or translation into a foreign language) without prior agreement and written consent from Agilent Technologies, Inc. as governed by United States and international copyright laws.

Agilent Technologies, Inc.
5301 Stevens Creek Blvd.
Santa Clara, CA 95051,
USA

Software Revision

This guide is valid for revision 2.7 of Agilent OpenLab CDS.

Warranty

The material contained in this document is provided "as is," and is subject to being changed, without notice, in future editions. Further, to the maximum extent permitted by applicable law, Agilent disclaims all warranties, either express or implied, with regard to this manual and any information contained herein, including but not limited to the implied warranties of merchantability and fitness for a particular purpose. Agilent shall not be liable for errors or for incidental or consequential damages in connection with the furnishing, use, or performance of this document or of any information contained herein. Should Agilent and the user have a separate written agreement with warranty terms covering the material in this document that conflict with these terms, the warranty terms in the separate agreement shall control.

Technology Licenses

The hardware and/or software described in this document are furnished under a license and may be used or copied only in accordance with the terms of such license.

Restricted Rights Legend

U.S. Government Restricted Rights. Software and technical data rights granted to the federal government include only those rights customarily provided to end user customers. Agilent provides this customary commercial license in Software and technical data pursuant to FAR 12.211 (Technical Data) and 12.212 (Computer Software) and, for the Department of Defense, DFARS 252.227-7015 (Technical Data - Commercial Items) and DFARS 227.7202-3 (Rights in Commercial Computer Software or Computer Software Documentation).

Safety Notices

CAUTION

A **CAUTION** notice denotes a hazard. It calls attention to an operating procedure, practice, or the like that, if not correctly performed or adhered to, could result in damage to the product or loss of important data. Do not proceed beyond a **CAUTION** notice until the indicated conditions are fully understood and met.

WARNING

A **WARNING** notice denotes a hazard. It calls attention to an operating procedure, practice, or the like that, if not correctly performed or adhered to, could result in personal injury or death. Do not proceed beyond a **WARNING** notice until the indicated conditions are fully understood and met.

In This Guide...

This guide addresses the advanced users, system administrators and persons responsible for validating Agilent OpenLab CDS. It contains reference information on the principles of calculations and data analysis algorithms.

Use this guide to verify system functionality against your user requirements specifications and to define and execute the system validation tasks defined in your validation plan. The following resources contain additional information.

- For context-specific task ("How To") information, references to the User Interface, and troubleshooting help: OpenLab Help and Learning.
- For details on system installation and site preparation: The *OpenLab CDS Requirements Guide*, *OpenLab CDS Workstation Guide* or *OpenLab CDS Client and AIC Guide*.

1 Signal Preparation

This chapter describes how the signal can be prepared, for example by blank subtraction, before it is integrated.

2 Integration with ChemStation Integrator

This chapter describes the concepts and integrator algorithms of the ChemStation integrator in OpenLab CDS.

3 Integration with EZChrom Integrator

This chapter contains the description of EZChrom integration events.

4 Peak Identification

This chapter describes the concepts of peak identification.

5 Calibration

This chapter contains details of the calculations used in the calibration process.

6 Quantitation

This chapter describes how compounds are quantified, and explains the calculations used in quantitation.

7 UV Spectral Analysis

This chapter describes the concepts of the impurity check and the confirmation of compound identity based on UV spectral analysis.

8 Mass Spectrometry

This chapter describes the sample purity calculation based on mass spectrometry.

9 System Suitability

This chapter describes what OpenLab CDS can do to evaluate the performance of both the analytical instrument and the analytical method.

Contents

1	Signal Preparation	7
	Signal Smoothing	8
	Blank Subtraction	11
2	Integration with ChemStation Integrator	12
	What is Integration?	13
	The Integrator Algorithms	15
	Principle of Operation	20
	Peak Recognition	21
	Peak Area Measurement	32
	Baseline Allocation	35
	Integration Events	46
3	Integration with EZChrom Integrator	66
	Integration Events	67
	Baseline Code Descriptions	83
4	Peak Identification	85
	What is Peak Identification?	86
	Conflict Resolution	88
	Relative Retention Times	89
	Time Reference Compound	91
	Update Processing Method	93
5	Calibration	98
	What is Calibration?	99
	Calibration Curve	100
	Calibration Curve Calculation	113
	Evaluating the Calibration Curve	121

6 Quantitation 127

- What is Quantitation? 128
- Correction Factors 129
- Concentration and Mass% 130
- Area% and Height% 131
- Quantitation of Calibrated Compounds 132
- Quantitation of Uncalibrated Compounds 137
- Quantitation of Not Identified Peaks 140
- Normalization 141
- Quantitation of groups 144

7 UV Spectral Analysis 152

- What is UV spectral analysis? 153
- UV impurity check 155
- UV confirmation 166

8 Mass Spectrometry 167

- MS sample purity 168
- MS peak purity 170

9 System Suitability 172

- Evaluating system suitability 173
- Noise Determination 175
- Calculation of peak asymmetry and symmetry 186
- System Suitability Formulas and Calculations 188
- Performance Test Definitions 189

Signal Smoothing	8
General approach	8
Algorithm details	9
Blank Subtraction	11

This chapter describes how the signal can be prepared, for example by blank subtraction, before it is integrated.

This chapter describes how the signal can be prepared before it is integrated.

NOTE

When both blank subtraction and smoothing settings are used to process data, the system will do smoothing first and then do blank subtraction with the smoothed signals.

Signal Smoothing

General approach

Assumptions

All smoothing algorithms assume that the data is equidistant data.

Non-equidistant data is transformed into equidistant data by applying a spline interpolation and resampling the data using the smallest time difference in the non-equidistant data.

Smoothing - base algorithm

All smoothing algorithms apply a window of size $2m+1$ populated with smoothing coefficients, using the following approach:

$$x'(i) = \sum_{j=-m}^{j=m} x(i+j) \cdot a(j)$$

where

a	Array of smoothing coefficients
x'	Smoothed signal
m	Even number specifying the half width of the smoothing window

This approach leads to an odd number $2m+1$ for the total window size.

Handling the edges

Since the smoothing coefficients are supposed to be normalized, the edges need special consideration.

For **Moving average** and **Gaussian** filtering, the window is pruned at the left or right edge, and the coefficients are recomputed to have a total sum of 1 (normalization).

For **Savitzky-Golay** the handling is more complicated. It needs to preserve the properties of Savitzky-Golay filtering also close to the edges of the signal. See ["Algorithm details"](#) on page 9.

Algorithm details

Moving average

The **Moving average** is the simplest smoothing algorithm. All coefficients are computed as follows:

$$a(i) = \frac{1}{2m+1}; i = [-m, m]$$

where

a	Array of smoothing coefficients
m	Even number specifying the half width of the smoothing window

This means the smoothing function has a rectangular shape. This kind of smoothing is also called boxcar averaging.

Gaussian

Gaussian smoothing uses coefficients sampled from the Gauss normal distribution. Given the number $2m+1$ as window size, the standard deviation σ of the normal distribution is computed as follows:

$$\sigma = \frac{2m+1}{6} - 1$$

The individual coefficients are then computed as:

$$a'(i) = e^{-\frac{1}{2} \left(\frac{i}{\sigma} \right)^2}, i = [-m, m]$$

In a second step the coefficients are normalized such that the sum is 1:

$$a(i) = \frac{a'(i)}{\sum_{j=-m}^m a'(j)}$$

Savitzky-Golay

The coefficients for Savitzky-Golay smoothing are computed in a way that ensures that the area under the function remains unchanged. The computation of coefficients is based on the paper *General Least-Squares Smoothing and Differentiation by the Convolution (Savitzky-Golay) Method* by Gorry (1990)¹.

This computation makes sure that the area-preserving properties of the Savitzky-Golay smoothing are also valid at the borders of the signal.

¹ Gorry, P.A., 1990. General Least-Squares Smoothing and Differentiation by the Convolution (Savitzky-Golay) Method. Analytical Chemistry 62, 570-573.

Blank Subtraction

When analyzing a sample, the obtained signal may be caused by analytes as well as by dilution solvents, mobile phases, additives etc. Use the blank subtraction to receive a clean chromatogram with contribution of the analytes only.

Blank signals can origin from:

- a blank sample within a sequence
- a blank sample outside of the sequence (for example, a single run)

The new signal is calculated by subtracting the blank signal:

$$\text{New signal} = \text{sample signal} - \text{blank signal}$$

For extracted chromatograms using a specific wavelength:

$$\text{New signal} = \text{sample signal at wavelength} - \text{blank signal at wavelength}$$

If a blank and a sample have different data rates, the data rate of the blank is adjusted. Data points are removed or created by spline interpolation.

If the run time of the sample is longer than the run time of the blank, the new signal will contain corrected and not corrected data points.

What is Integration?	13
What Does Integration Do?	13
Integrator Capabilities	14
The Integrator Algorithms	15
Overview	15
Defining the Initial Baseline	16
Tracking the Baseline	17
Allocating the Baseline	18
Definition of Terms	19
Principle of Operation	20
Peak Recognition	21
Peak Width	21
Peak Recognition Filters	22
Bunching	23
The Peak Recognition Algorithm	24
Merged Peaks	26
Shoulders	27
Default Baseline Construction	28
Baseline Codes	29
Peak Area Measurement	32
Determination of the area	33
Units and Conversion Factors	34
Baseline Allocation	35
Baseline Correction Modes	35
Peak-to-Valley Ratio	37
Tangent Skimming	38
Tangent Skim Modes	42
Integration Events	46
Standard Integration Events: Initial Events	46
Standard Integration Events: Timed Events	51
Advanced Integration Events	63

This chapter describes the concepts and integrator algorithms of the ChemStation integrator in OpenLab CDS.

What is Integration?

Integration locates the peaks in a signal and calculates their size.

Integration is a necessary step for:

- identification
- calibration
- quantitation

What Does Integration Do?

When a signal is integrated, the software:

- identifies a start and an end time for each peak
- finds the apex of each peak; that is, the retention/migration time,
- constructs a baseline, and
- calculates the area, height, peak width, and symmetry for each peak.

This process is controlled by parameters called integration events.

Integrator Capabilities

The integrator algorithms include the following key capabilities:

- the ability to define individual integration event tables for each chromatographic signal if multiple signals or more than one detector are used
- graphical manual integration of chromatograms requiring human interpretation
- annotation of integration results
- integrator parameter definitions to set or modify the basic integrator settings for area rejection, height rejection, peak width, slope sensitivity, shoulder detection, baseline correction and front/tail tangent skim detection
- baseline control parameters, such as force baseline, hold baseline, baseline at all valleys, baseline at the next valley, fit baseline backwards from the end of the current peak, most likely baseline point from a time range
- area summation control
- negative peak recognition
- solvent peak definition detection
- integrator control commands defining retention time ranges for the integrator operation
- peak shoulder allocation through the use of second derivative calculations

The Integrator Algorithms

Overview

To integrate a chromatogram, the integrator ...

- 1 defines the initial baseline,
- 2 continuously tracks and updates the baseline,
- 3 identifies the start time for a peak,
- 4 finds the apex of each peak,
- 5 identifies the end time for the peak,
- 6 constructs a baseline, and
- 7 calculates the area, height, peak width, and symmetry for each peak.

This process is controlled by **integration events**. The most important events are initial slope sensitivity, peak width, shoulders mode, area reject, and height reject. The software allows you to set initial values for these and other events. The initial values take effect at the beginning of the chromatogram.

In most cases, the initial events will give good integration results for the entire chromatogram, but there may be times when you want more control over the progress of an integration.

The software allows you to control how an integration is performed by enabling you to program new integration events at appropriate times in the chromatogram.

Defining the Initial Baseline

Cardinal points

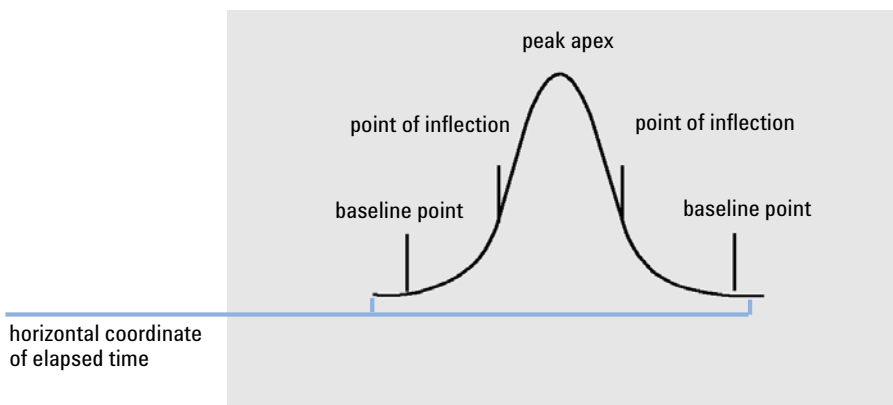


Figure 1 Cardinal points

Defining the initial baseline

Because baseline conditions vary according to the application and detector hardware, the integrator uses parameters from both the integration events and the data file to optimize the baseline.

Before the integrator can integrate peaks, it must establish a **baseline point**. At the beginning of the analysis, the integrator establishes an initial baseline level by taking the first data point as a tentative baseline point. It then attempts to redefine this initial baseline point based on the average of the input signal. If the integrator does not obtain a redefined initial baseline point, it retains the first data point as a potential initial baseline point.

Identifying the cardinal points of a peak

The integrator determines that a peak may be starting when potential baseline points lie outside the baseline envelope, and the baseline curvature exceeds a certain value, as determined by the integrator's slope sensitivity parameter. If this condition continues, the integrator recognizes that it is on the up-slope of a peak, and the peak is processed.

Start

- 1 Slope and curvature within limit: continue tracking the baseline.
- 2 Slope and curvature above limit: possibility of a peak.
- 3 Slope remains above limit: peak recognized, peak start point defined.
- 4 Curvature becomes negative: front inflection point defined.

Apex

- 1 Slope passes through zero and becomes negative: apex of peak, apex point defined.
- 2 Curvature becomes positive: rear inflection point defined.

End

- 1 Slope and curvature within limit: approaching end of the peak.
- 2 Slope and curvature remain within limit: end of peak defined.
- 3 The integrator returns to the baseline tracking mode.

Tracking the Baseline

The integrator samples the digital data at a rate determined by the initial peak width or by the calculated peak width, as the run progresses. It considers each data point as a potential baseline point.

The integrator determines a *baseline envelope* from the slope of the baseline, using a baseline-tracking algorithm in which the slope is determined by the first derivative and the curvature by the second derivative. The baseline envelope can be visualized as a cone, with its tip at the current data point. The upper and lower acceptance levels of the cone are:

- $+ \text{upslope} + \text{curvature} + \text{baseline bias}$ must be lower than the threshold level,
- $- \text{upslope} - \text{curvature} + \text{baseline bias}$ must be more positive (i.e. less negative) than the threshold level.

As new data points are accepted, the cone moves forward until a break-out occurs.

To be accepted as a baseline point, a data point must satisfy the following conditions:

- it must lie within the defined baseline envelope,
- the curvature of the baseline at the data point (determined by the derivative filters), must be below a critical value, as determined by the current slope sensitivity setting.

The initial baseline point, established at the start of the analysis is then continuously reset, at a rate determined by the peak width, to the moving average of the data points that lie within the baseline envelope over a period determined by the peak width. The integrator tracks and periodically resets the baseline to compensate for drift, until a peak up-slope is detected.

Allocating the Baseline

The integrator allocates the chromatographic baseline during the analysis at a frequency determined by the peak width value. When the integrator has sampled a certain number of data points, it resets the baseline from the initial baseline point to the current baseline point. The integrator resumes tracking the baseline over the next set of data points and resets the baseline again. This process continues until the integrator identifies the start of a peak.

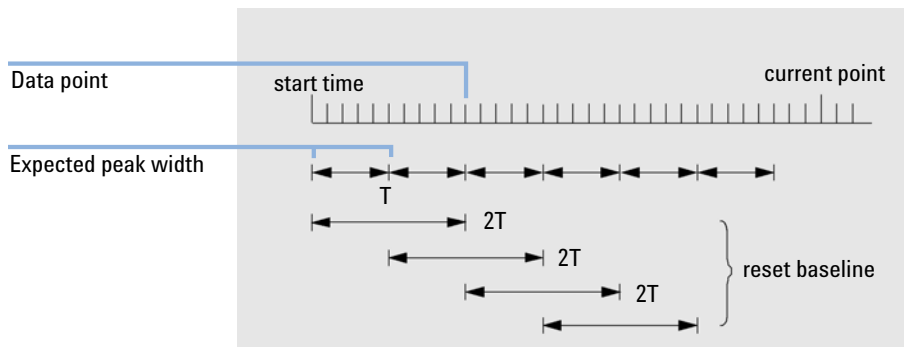


Figure 2 Baseline

At the start of the integration process the first data point is used. This baseline point is periodically reset as shown in the figure (see [Figure 2](#) on page 18).

Areas are summed over a time T (expected peak width). This time can never be shorter than one data point. This continues as long as baseline condition exists. Slope and curvature are also taken. If both slope and curvature are less than the threshold, two summed areas are added together, and compared with the

previous baseline. If the new value is less than the previous baseline, the new value immediately replaces the old one. If the new value is greater than the previous value, it is stored as a tentative new baseline value and is confirmed if one more value satisfies slope and curvature flatness criteria. This latter limitation is not in effect if negative peaks are allowed. During baseline, a check must also be made to examine fast rising solvents. They may be too fast for upslope detection. (By the time upslope is confirmed, solvent criterion may no longer be valid.) At first time through the first data point is baseline. It is replaced by the 2 T average if signal is on base. Baseline is then reset every T (see [Figure 2](#) on page 18).

Definition of Terms

Solvent peak

The solvent peak, which is generally a very large peak of no analytical importance, is not normally integrated. However, when small peaks of analytical interest elute close to the solvent peak, for example, on the tail of the solvent peak, special integration conditions can be set up to calculate their areas corrected for the contribution of the solvent peak tail.

Shoulder (front, rear)

Shoulders occur when two peaks elute so close together that no valley exists between them, and they are unresolved. Shoulders may occur on the leading edge (front) of the peak, or on the trailing edge (rear) of the peak. When shoulders are detected, they may be integrated either by tangent skim or by drop-lines.

Slope

The slope of a peak, which denotes the change of concentration of the component against time, is used to determine the onset of a peak, the peak apex, and the end of the peak.

Principle of Operation

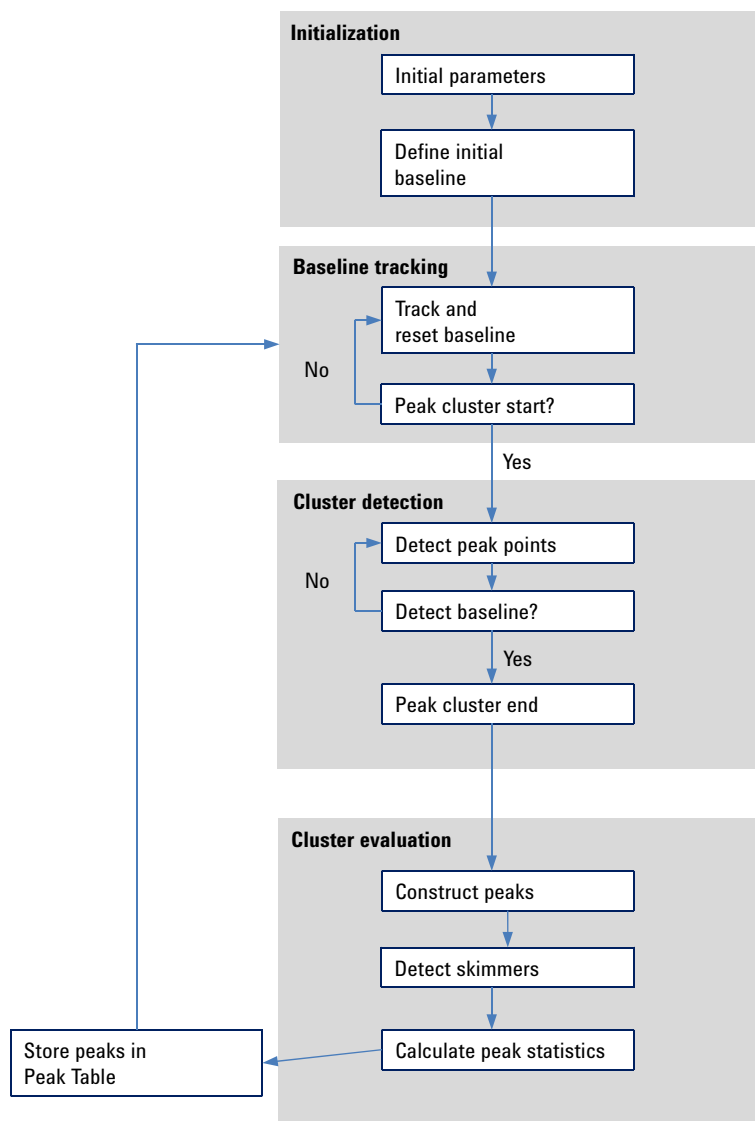


Figure 3 Integrator Flow Diagram

Peak Recognition

The integrator uses several tools to recognize and characterize a peak:

- "Peak Width" on page 21
- "Peak Recognition Filters" on page 22
- "Bunching" on page 23
- "The Peak Recognition Algorithm" on page 24
- "Merged Peaks" on page 26
- "Shoulders" on page 27
- "Default Baseline Construction" on page 28
- "Baseline Codes" on page 29

Peak Width

During integration, the peak width is calculated from the adjusted peak area and height:

$$\text{Width} = \text{adjusted area} / \text{adjusted height}$$

or, if the inflection points are available, from the width between the inflection points.

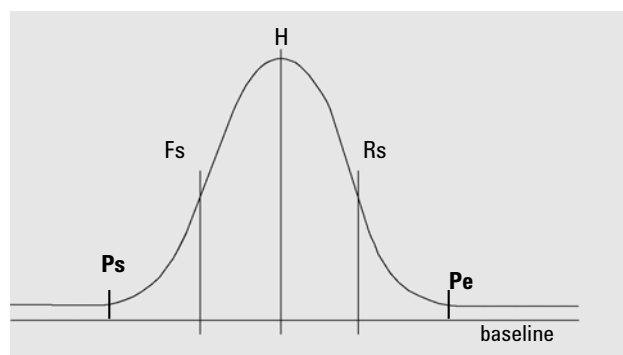


Figure 4 Peak width calculation

In the figure above, the total area, A , is the sum of the areas from peak start (P_s) to peak end (P_e), adjusted for the baseline. F_s is the front slope at the inflection point, R_s is the rear slope at the inflection point.

The peak width setting controls the ability of the integrator to distinguish peaks from baseline noise. To obtain good performance, the peak width must be set close to the width of the actual chromatographic peaks.

There are three ways the peak width is changed:

- before the integration process, you can specify the initial peak width,
- during the integration process, the integrator automatically updates the peak width as necessary to maintain a good match with the peak recognition filters,
- during the integration process, you can reset or modify the peak width using a time-programmed event.

Peak Recognition Filters

The integrator has three peak recognition filters that it can use to recognize peaks by detecting changes in the slope and curvature within a set of contiguous data points. These filters contain the first derivative (to measure slope) and the second derivative (to measure curvature) of the data points being examined by the integrator.

NOTE

In order to deliver reliable results the peak must contain at least ten data points.

The recognition filters are:

- Filter 1** Slope (curvature) of two (three) contiguous data points
- Filter 2** Slope of four contiguous data points and curvature of three non-contiguous data points
- Filter 3** Slope of eight contiguous data points and curvature of three non-contiguous data points

The actual filter used is determined by the peak width setting. For example, at the start of an analysis, Filter 1 may be used. If the peak width increases during the analysis, the filter is changed first to Filter 2 and then to Filter 3. To obtain good performance from the recognition filters, the peak width must be set close to the width of the actual chromatographic peaks. During the run, the integrator updates the peak width as necessary to optimize the integration.

The integrator calculates the updated peak width in different ways, depending on the instrument technique.

For LC data, the default peak width calculation uses a composite calculation:

$$0.3 * (\text{Right Inflection Point} - \text{Left Inflection Point}) + 0.7 * \text{Area} / \text{Height}$$

For GC data, the default peak width calculation uses area/height. This calculation does not overestimate the width when peaks are merged above the half-height point.

In certain types of analysis, for example isothermal GC and isocratic LC analyses, peaks become significantly broader as the analysis progresses. To compensate for this, the integrator automatically updates the peak width as the peaks broaden during the analysis. It does this automatically unless the updating has been disabled with the fixed peak width timed event.

The peak width update is weighted in the following way:

$$0.75 * (\text{Existing Peak Width}) + 0.25 * (\text{Width of Current Peak})$$

Bunching

Bunching is the means by which the integrator keeps broadening peaks within the effective range of the peak recognition filters to maintain good selectivity.

The integrator cannot continue indefinitely to increase the peak width for broadening peaks. Eventually, the peaks would become so broad that they could not be seen by the peak recognition filters. To overcome this limitation, the integrator bunches the data points together, effectively narrowing the peak while maintaining the same area.

When data is bunched, the data points are bunched as two raised to the bunching power, i.e. unbunched = 1x, bunched once = 2x, bunched twice = 4x etc.

Bunching is based on the data rate and the peak width. The integrator uses these parameters to set the bunching factor to give the appropriate number of data points (see [Table 1](#) on page 24).

Bunching is performed in the powers of two based on the expected or experienced peak width. The bunching algorithm is summarized in [Table 1](#) on page 24.

Table 1 Bunching criteria

Expected Peak Width	Filter(s) Used	Bunching Done
0 - 10 data points	First	None
8 - 16 data points	Second	None
12 - 24 data points	Third	None
16 - 32 data points	Second	Once
24 - 48 data points	Third	Once
32 - 96 data points	Third, second	Twice
64 - 192 data points	Third, second	Three times

The Peak Recognition Algorithm

The integrator identifies the start of the peak with a baseline point determined by the peak recognition algorithm. The peak recognition algorithm first compares the outputs of the peak recognition filters with the value of the initial slope sensitivity, to increase or decrease the up-slope accumulator. The integrator declares the point at which the value of the up-slope accumulator is ≥ 15 the point that indicates that a peak has begun.

Peak Start

In [Table 2](#) on page 25 the expected peak width determines which filter's slope and curvature values are compared with the Slope Sensitivity. For example, when the expected peak width is small, Filter 1 numbers are added to the up-slope accumulator. If the expected peak width increases, then the numbers for Filter 2 and, eventually, Filter 3 are used.

When the value of the up-slope accumulator is ≥ 15 , the algorithm recognizes that a peak may be starting.

Table 2 Incremental Values to Upslope Accumulator

Derivative Filter 1 - 3 Outputs against Slope Sensitivity	Filter 1	Filter 2	Filter 3
Slope > Slope Sensitivity	+8	+5	+3
Curvature > Slope Sensitivity	+0	+2	+1
Slope < (-) Slope Sensitivity	-8	-5	-3
Slope < Slope Sensitivity	-4	-2	-1
Curvature < (-) Slope Sensitivity	-0	-2	-1

Peak End

In [Table 3](#) on page 25 the expected peak width determines which filter's slope and curvature values are compared with the Slope Sensitivity. For example, when the expected peak width is small, Filter 1 numbers are added to the down-slope accumulator. If the expected peak width increases, then the numbers for Filter 2 and, eventually, Filter 3 are used.

When the value of the down-slope accumulator is ≥ 15 , the algorithm recognizes that a peak may be ending.

Table 3 Incremental Values for Downslope Accumulator

Derivative Filter 1 - 3 Outputs against Slope Sensitivity	Filter 1	Filter 2	Filter 3
Slope < (-) Slope Sensitivity	+8	+5	+3
Curvature < (-) Slope Sensitivity	+0	+2	+1
Slope > Slope Sensitivity	-11	-7	-4
Slope > Slope Sensitivity	-28	-18	-11
Curvature > Slope Sensitivity	-0	-2	-1

The Peak Apex Algorithm

The peak apex is recognized as the highest point in the chromatogram by constructing a parabolic fit that passes through the highest data points.

Merged Peaks

Merged peaks occur when a new peak begins before the end of peak is found. The figure illustrates how the integrator deals with merged peaks.

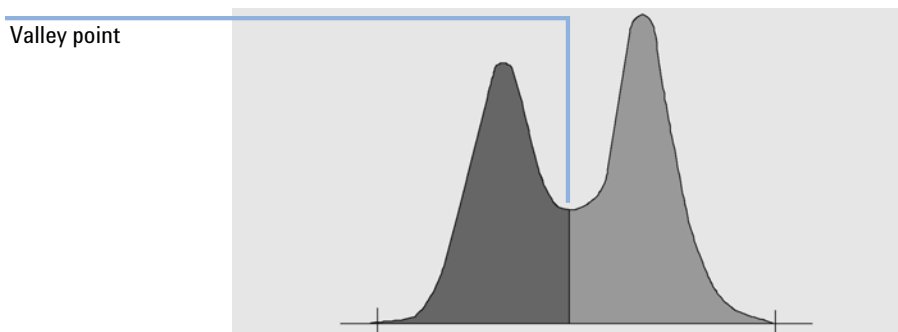


Figure 5 Merged Peaks

The integrator processes merged peaks in the following way:

- 1 it sums the area of the first peak until the valley point.
- 2 at the valley point, area summation for the first peak ends and summation for the second peak begins.
- 3 when the integrator locates the end of the second peak, the area summation stops. This process can be visualized as separating the merged peaks by dropping a perpendicular from the valley point between the two peaks.

Shoulders

Shoulders are unresolved peaks on the leading or trailing edge of a larger peak. When a shoulder is present, there is no true valley in the sense of negative slope followed by positive slope. A peak can have any number of front and/or rear shoulders.

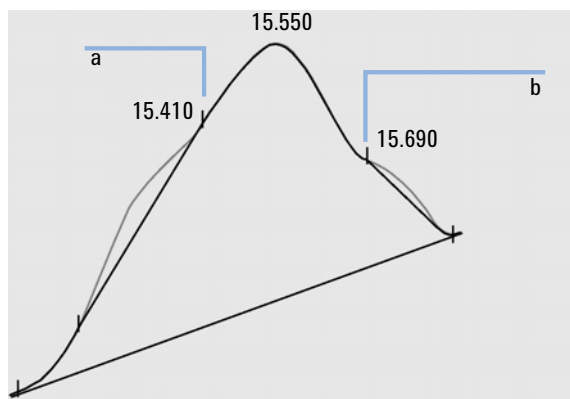


Figure 6 Peak Shoulders

Shoulders are detected from the curvature of the peak as given by the second derivative. When the curvature goes to zero, the integrator identifies a point of inflection, such as points a and b in Figure 6 on page 27.

- A potential front shoulder exists when a second inflection point is detected before the peak apex. If a shoulder is confirmed, the start of the shoulder point is set at the maximum positive curvature point before the point of inflection.
- A potential rear shoulder exists when a second inflection point is detected before the peak end or valley. If a shoulder is confirmed, the start of the shoulder point is set at the point of the first minimum of the slope after the peak apex.

Retention time is determined from the shoulder's point of maximum negative curvature. With a programmed integration event, the integrator can also calculate shoulder areas as normal peaks with drop-lines at the shoulder peak points of inflection.

The area of the shoulder is subtracted from the main peak.

Peak shoulders can be treated as normal peaks by use of an integrator timed event.

Default Baseline Construction

After any peak cluster is complete, and the baseline is found, the integrator requests the baseline allocation algorithm to allocate the baseline using a pegs-and-thread technique. It uses trapezoidal area and proportional height corrections to normalize and maintain the lowest possible baseline. Inputs to the baseline allocation algorithm also include parameters from the method and data files that identify the detector and the application, which the integrator uses to optimize its calculations.

In the simplest case, the integrator constructs the baseline as a series of straight line segments between:

- the start of baseline,
- peakstart, valley, end points,
- the peak baseline

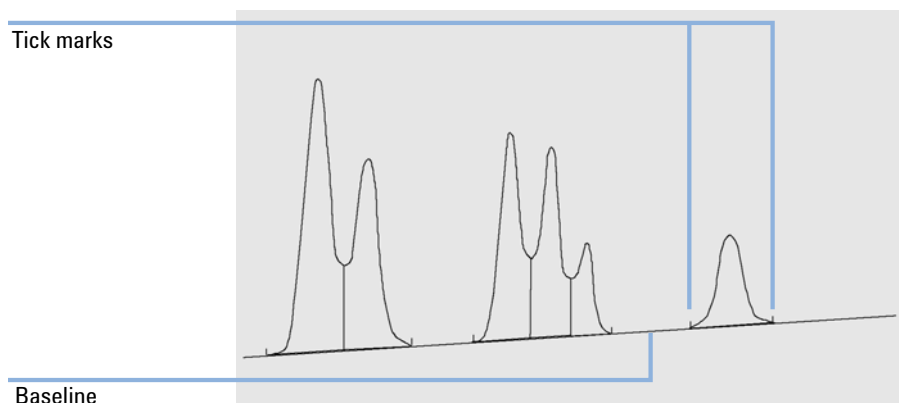


Figure 7 Default Baseline Construction

Baseline Codes

In the integration results of a report, each peak is assigned a two-, three- or four-character code that describes how the signal baseline was drawn.

Table 4 Four character code

First character	Second character	Third character	Fourth character
Baseline at start	Baseline at end	Error/peak flag	Peak type

The baseline codes are included in the **Injection Results** table and in all default report templates.

Injection Results								
Peaks		Summary						
#	Name	BL code	Signal description	RT (min)	Area	Area%	Height	
1		VB	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	1.940	356.576	27.285	65.775	
2		BB	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	2.578	314.315	24.051	49.125	
3		BB	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	3.022	333.322	25.506	59.034	
4		BV	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	4.517	302.650	23.159	49.419	

Figure 8 Injection results

Signal: DAD1A,Sig=272.0,16.0 Ref=380.0,20.0					
RT [min]	Type	Width [min]	Area	Height	Area%
1.940	VB	0.53	356.58	65.78	27.28
2.578	BB	0.43	314.31	49.12	24.05
3.022	BB	0.64	333.32	59.03	25.51
4.517	BV	0.59	302.65	49.42	23.16
		Sum	1306.86		

Figure 9 Example: Table from Short Area report

Characters 1 and 2

The first character describes the baseline at the start of the peak and the second character describes the baseline at the end of the peak.

- B** The peak started or stopped on the baseline.
- P** The peak started or stopped while the baseline was penetrated.
- V** The peak started or stopped with a valley drop-line.
- H** The peak started or stopped on a forced horizontal baseline.
- F** The peak started or stopped on a forced point.
- M** The peak was manually integrated.
- U** The peak was unassigned.

Additional flags may also be appended (in order of precedence).

Character 3

The third character describes an error or peak flag:

- A** The integration was aborted. For example due to the integration events ON/OFF, or due to the end of signal run time.
- D** The peak was distorted (bad peak shape).

Blank space The peak is a normal peak.

Character 4

The fourth character describes the peak type. It is shown only for forced integration events, or if manual integration has been triggered. For example, you use an integration event to define a solvent peak, or you use manual integration to correct the baseline or to delete a peak.

- S** The peak is a solvent peak.
- N** The peak is a negative peak.
- +** The peak is an area summed peak.
- T** Tangent-skimmed peak (standard skim).
- X** Tangent-skimmed peak (old mode exponential skim).
- E** Tangent-skimmed peak (new mode exponential skim).
- m** Peak defined by manual baseline.
- n** Negative peak defined by manual baseline.
- t** Tangent-skimmed peak defined by manual baseline.
- x** Tangent-skimmed peak (exponential skim) defined by manual baseline.
- R** The peak is a recalculated peak.
- f** Peak defined by a front shoulder tangent.
- b** Peak defined by a rear shoulder tangent.
- F** Peak defined by a front shoulder drop-line.
- B** Peak defined by a rear shoulder drop-line.
- U** The peak is unassigned.

Peak Area Measurement

The final step in peak integration is determining the final area of the peak.

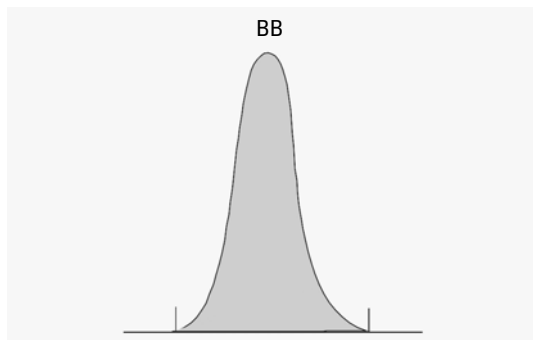


Figure 10 Area measurement for Baseline-to-Baseline Peaks

In the case of a simple, isolated peak, the peak area is determined by the accumulated area above the baseline between peak start and stop.

Determination of the area

The area that the integrator calculates during integration is determined as follows:

- for baseline-to-baseline (BB) peaks, the area above the baseline between the peak start and peak end, as in [Figure 10](#) on page 32,
- for valley-to-valley (VV) peaks, the area above the baseline, segmented with vertical dropped lines from the valley points, as in [Figure 11](#) on page 33,

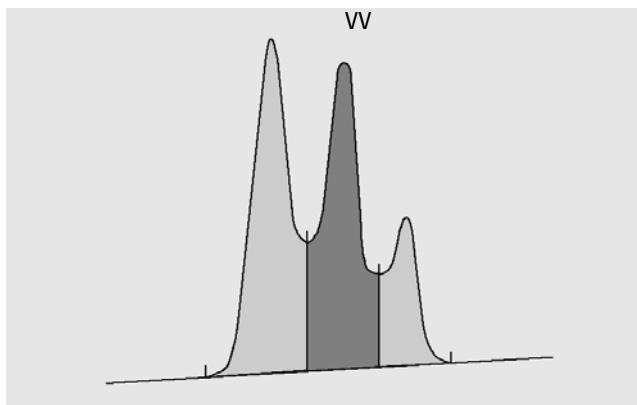


Figure 11 Area Measurement for Valley-to-Valley Peaks

- for tangent (T) peaks, the area above the reset baseline,
- for solvent (S) peaks, the area above the horizontal extension from the last-found baseline point and below the reset baseline given to tangent (T) peaks. A solvent peak may rise too slowly to be recognized, or there may be a group of peaks well into the run which you feel should be treated as a solvent with a set of riders. This usually involves a merged group of peaks where the first one is far larger than the rest. The simple drop-line treatment would exaggerate the later peaks because they are actually sitting on the tail of the

first one. By forcing the first peak to be recognized as a solvent, the rest of the group is skimmed off the tail,

- negative peaks that occur below the baseline have a positive area, as shown in Figure 12 on page 34.

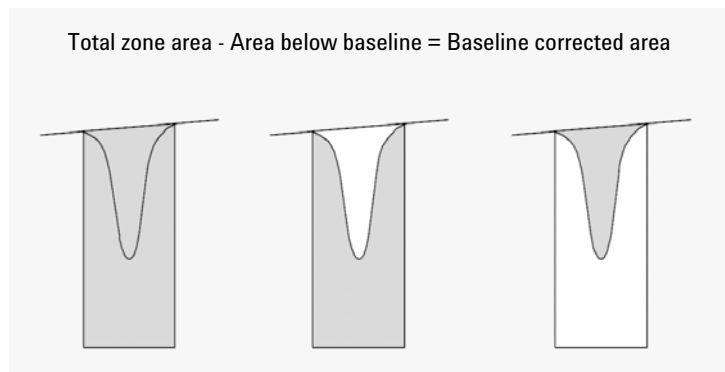


Figure 12 Area Measurement for Negative Peaks

Units and Conversion Factors

Externally, the data contains a set of data points; they can be either sampled data or integrated data. In the case of integrated data, each data point corresponds to an area, which is expressed as *Height × Time*. In the case of sampled data, each data point corresponds to a height.

Therefore, in the case of integrated data, height is a calculated entity, obtained by dividing area by the time elapsed since the preceding data point. In the case of sampled data, area is calculated by multiplying the data by the time elapsed since the preceding data point.

The integration calculation makes use of both entities. The units carried internally inside the integrator are: *detector response × seconds* for area, and **detector response** as height. This is done to provide a common base for integer truncations when needed. The measurements of time, area and height are reported in real physical units, irrespective of how they are measured, calculated and stored in the software.

Baseline Allocation

Baseline Correction Modes

In OpenLab CDS, several baseline correction modes are available. They are described in the following sections.

Baseline Correction Mode: Classical

A penetration occurs when the signal drops below the constructed baseline (point a in [Figure 13](#) on page 35).

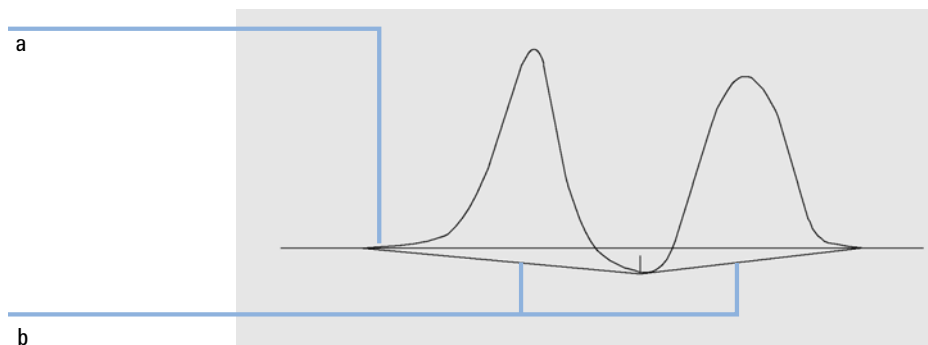


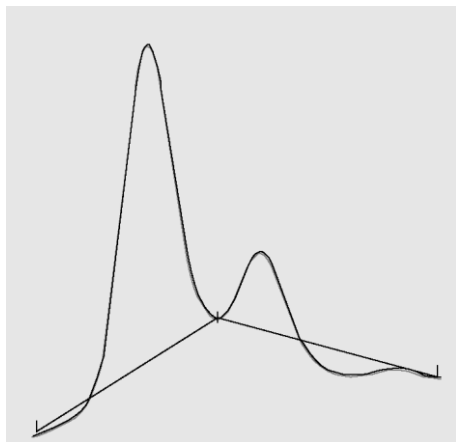
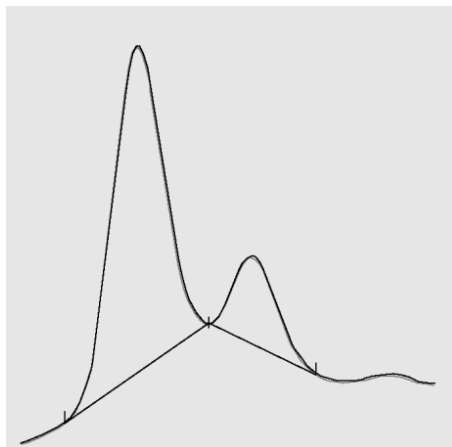
Figure 13 Baseline Penetration

If a baseline penetration occurs, that part of the baseline may be reconstructed, as shown by points b in [Figure 13](#) on page 35. You can use the following correction modes to remove all baseline penetrations:

- **No penetration**
- **Advanced**

Baseline Correction Mode: No Penetration

When this option is selected, each peak cluster is searched for baseline penetrations. If penetrations are found, the start and/or end points of the peak are shifted until there are no penetrations left.

Baseline correction mode **Classical**Baseline correction mode **No penetration****NOTE**

The baseline correction mode **No penetration** is not available for solvent peaks, with their child peaks and shoulders.

Baseline Correction Mode: Advanced

In the advanced baseline correction mode, the integrator tries to optimize the start and end locations of the peaks, re-establishes the baseline for a cluster of peaks, and removes baseline penetrations (see [“Baseline Correction Mode: No Penetration”](#) on page 36). In many cases, advanced baseline correction gives a more stable baseline, which is less dependent on slope sensitivity.

Peak-to-Valley Ratio

The Peak to valley ratio is a measure of quality, indicating how well the peak is separated from other substance peaks. This user-specified parameter is a constituent of advanced baseline tracking mode. It is used to decide whether two peaks that do not show baseline separation are separated using a drop line or a valley baseline. The integrator calculates the ratio between the baseline-corrected height of the smaller peak and the baseline-corrected height of the valley. When the peak valley ratio is lower than the user-specified value, a drop-line is used; otherwise, a baseline is drawn from the baseline at the start of the first peak to the valley, and from the valley to the baseline at the end of the second peak (compare “Baseline Correction Mode: No Penetration” on page 36 with Figure 14 on page 37).

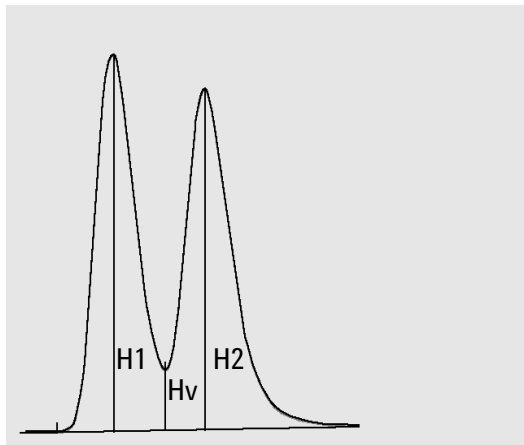


Figure 14 Peak Valley Ratio

The peak-to-valley ratio is calculated using the following equations:

$$H1 \geq H2, \text{ Peak valley ratio} = H2/H_v$$

and

$$H1 < H2, \text{ Peak valley ratio} = H1/H_v$$

Figure 15 on page 38 shows how the user-specified value of the peak valley ratio affects the baselines.

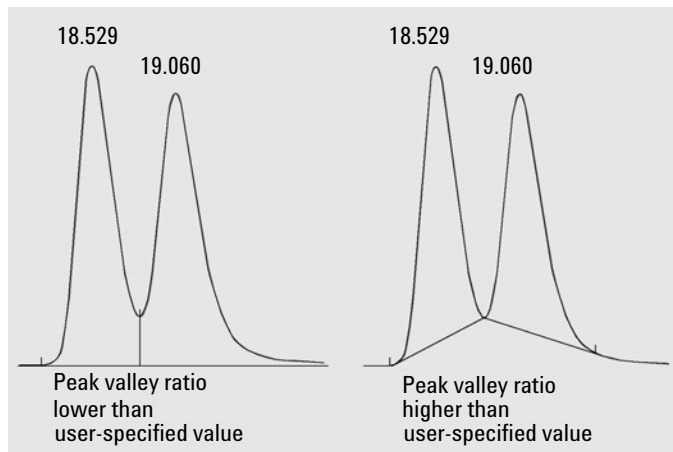


Figure 15 Effect of peak valley ratio on the baselines

Tangent Skimming

Tangent skimming is a form of baseline constructed for peaks found on the upslope or downslope of a peak. The prerequisite is that the two peaks are not baseline-separated.

The following figures illustrate the principle of tangent skimming:

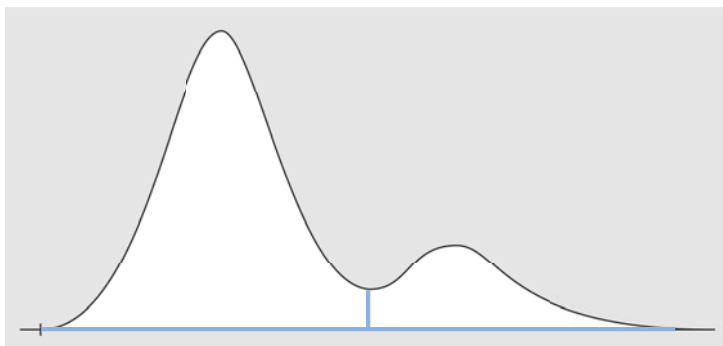


Figure 16 Peaks without skimming, separated by a drop line

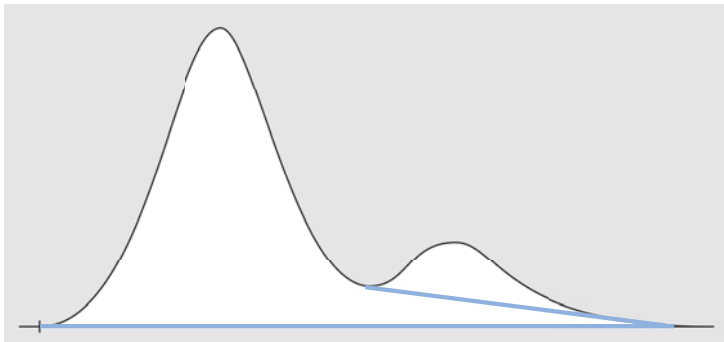


Figure 17 Tail skimming

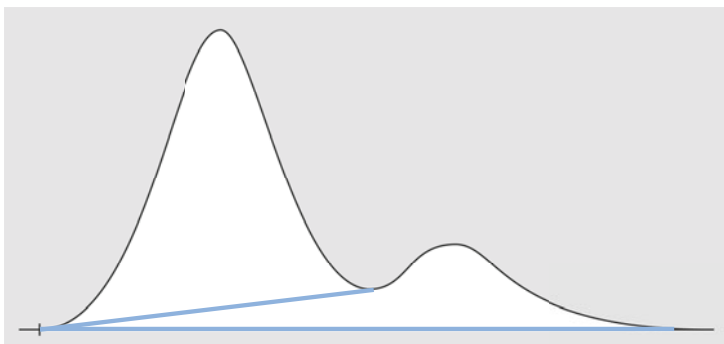


Figure 18 Front skimming

Skim Criteria

The following criteria determine whether a skim line is used to calculate the area of a child peak eluting on the leading or trailing edge of a parent peak:

- Skim height ratio (**Front skim height ratio** or **Tail skim height ratio**)
- **Skim valley ratio**

The *skim height ratio* is the ratio of the baseline-corrected height of the parent peak (H_p in the figure below) to the baseline-corrected height of the child peak (H_c). To have the child peak skimmed, use a value lower than this ratio. To disable exponential skimming throughout a run, you can set this parameter to a high value or to zero.

The *skim valley ratio* is the ratio of the height of the child peak above the baseline (H_c in the figure below) to the height of the valley above the baseline (H_v). To have the child peak skimmed, use a value greater than this ratio.

NOTE

If one of these criteria is not met for a set of child peaks at the tail of the parent peak, all child peaks after the last child peak that met both criteria are not skimmed anymore but use a drop line

NOTE

These criteria are not used if a timed event for an exponential is in effect, or if the parent peak is itself a child peak. The baseline code between parent peak and child peak must be of type **Valley** (see "Baseline Codes" on page 29).

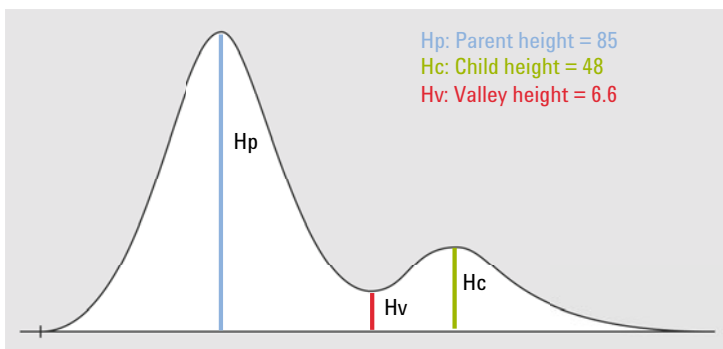


Figure 19 Example for calculating the skim criteria values

$$\text{Skim height ratio} = H_p / H_c$$

$$\text{Skim valley ratio} = H_c / H_v$$

where

H_p	Baseline-corrected height of parent peak
H_v	Height of valley above the baseline
H_c	Baseline-corrected height of child peak

Tail Skimming To use tail skimming, you would set the parameters as follows:

- Tail skim height ratio = $85 / 48 = 1.77$
In the integration events, use a value < 1.77 .
- Skim valley ratio = $48 / 6.6 = 7.3$
In the integration events, use a value > 7.3 .

Front Skimming

With front skimming, the first peak is the child peak, and the second peak is the parent peak. Thus, to use front skimming, you would set the parameters as follows:

- Front skim height ratio = $48 / 85 = 0.56$
In the integration events, use a value < 0.56 .
- Skim valley ratio = $85 / 6.6 = 12.9$
In the integration events, use a value > 12.9 .

Tangent Skim Modes

When tangent skimming is enabled, four models are available to calculate suitable peak areas:

- Exponential curve
- New exponential skim
- Straight line skim
- Combined exponential and straight line calculations for the best fit (standard skims)

Exponential Curve

This skim model draws a curve using an exponential equation through the start and end of the child peak. The curve passes under each child peak that follows the parent peak; the area under the skim curve is subtracted from the child peaks and added to the parent peak.

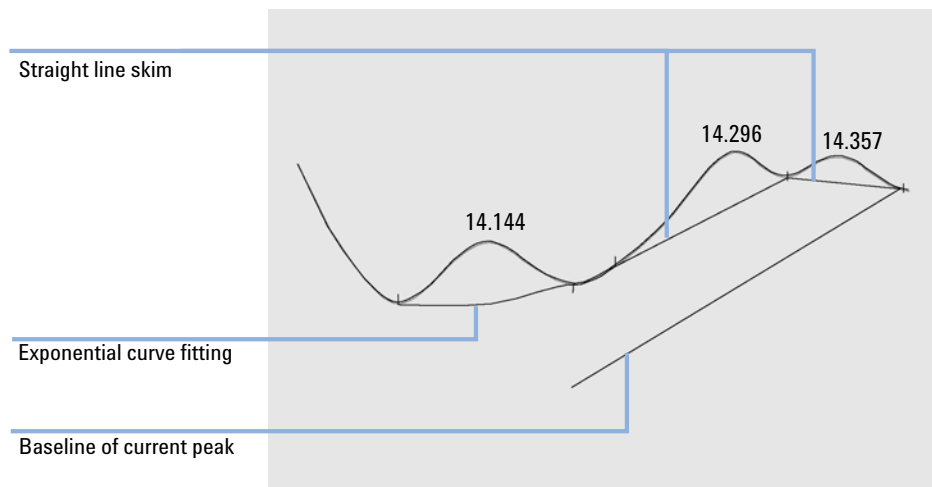
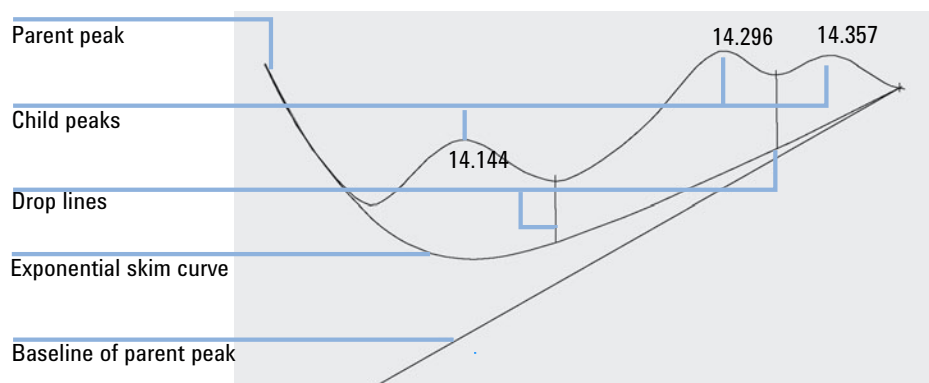


Figure 20 Exponential skim

New exponential curve

This skim model draws a curve using an exponential equation to approximate the leading or trailing edge of the parent peak. The curve passes under one or more peaks that follow the parent peak (child peaks). The area under the skim curve is subtracted from the child peaks and added to the main peak. More than one child peak can be skimmed using the same exponential model; all peaks after the first child peak are separated by drop lines, beginning at the end of the first child peak, and are dropped only to the skim curve.

**Figure 21** New exponential skim

Straight Line Skim

This skim model draws a straight line through the start and end of a child peak. The height of the start of the child peak is corrected for the parent peak slope. The area under the straight line is subtracted from the child peak and added to the parent peak.

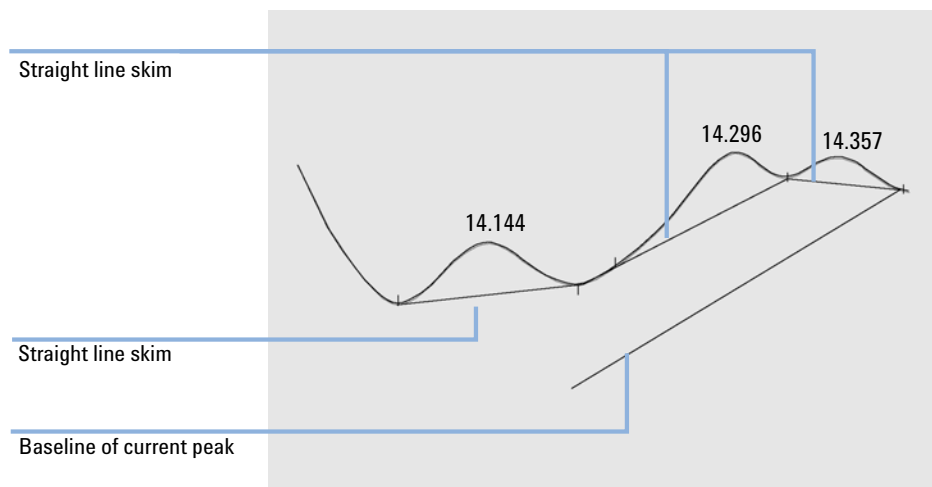


Figure 22 Straight line skim

Standard Skims

This default method is a combination of exponential and straight line calculations for the best fit.

The switch from an exponential to a linear calculation is performed in a way that eliminates abrupt discontinuities of heights or areas.

- When the signal is well above the baseline, the tail-fitting calculation is exponential.
- When the signal is within the baseline envelope, the tail fitting calculation is a straight line.

The combination calculations are reported as exponential or straight tangent skim.

Calculation of Exponential Curve for Skims

The following equation is used to calculate an exponential skim:

$$H_b(t_R) = H_0 * \exp(-B * (t_R - t_0)) + A * t_R + C$$

where

H_b	Height of the exponential skim at time t_R
H_0	Height (above baseline) of the start of the exponential skim
B	Decay factor of the exponential function
t_0	Time corresponding to the start of the exponential skim
t_R	Retention time
A	Slope of the baseline of the parent peak
C	Offset of the baseline of the parent peak

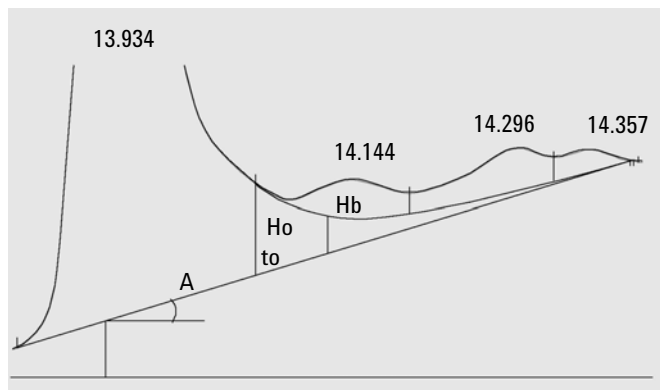


Figure 23 Values used to calculate an exponential skim

Integration Events

The available integration events are divided into the following groups:

- Initial integration events are those that apply from the start of integration. You can find them as default values in the **Standard** node of the **Integration Events** section in the processing method. These events cannot be deleted, but you may change the values.
- Timed events take place after the start of integration. Timed events may change the value of an initial event, or may switch on or off additional integration parameters. They can be added in the **Standard** node of the **Integration Events** section in the processing method.
- Integration events that always apply to all signals can be configured in the **Advanced** node of the **Integration Events** section.

Standard Integration Events: Initial Events

Slope sensitivity

Sets the value of the signal slope that is used to identify the start and end points of a peak during integration.

You can set the values either specifically for a given signal or globally for all signals.

When the signal slope exceeds the **Slope Sensitivity** value, a peak start point is established; when the signal slope decreases below the **Slope Sensitivity** value, a peak end point is established.

Peak width

Controls the selectivity of the integrator to distinguish peaks from baseline noise. You specify the peak width in units of time that correspond to the peak width at half-height of the first expected peak (excluding the solvent peak).

The integrator updates the peak width when necessary during the run to optimize the integration:

If the selected initial peak width is too low, noise may be interpreted as peaks. If broad and narrow peaks are mixed, you may decide to use runtime programmed events to adjust the peak width for certain peaks. Sometimes, peaks become significantly broader as the analysis progresses, for example in isothermal GC and isocratic LC analyses. To compensate for this, the integrator automatically updates the peak width as peaks broaden during an analysis unless disabled with a timed event.

The Peak Width update is weighted in the following way:

$$0,75 \times (\text{existing peak width}) + 0,25 \times (\text{width of current peak})$$

Area reject Sets the area of the smallest peak of interest.

Any peaks that have areas less than the minimum area are not reported: The integrator rejects any peaks that are smaller than the **Area Reject** value after baseline correction. The **Area Reject** value must be greater than or equal to zero.

NOTE

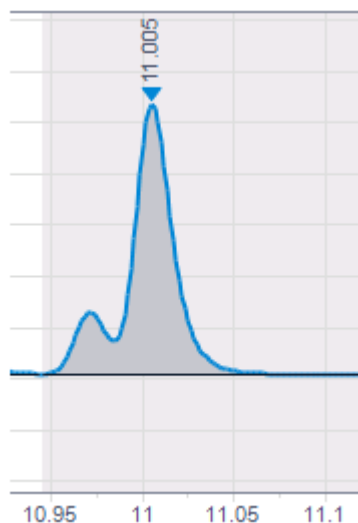
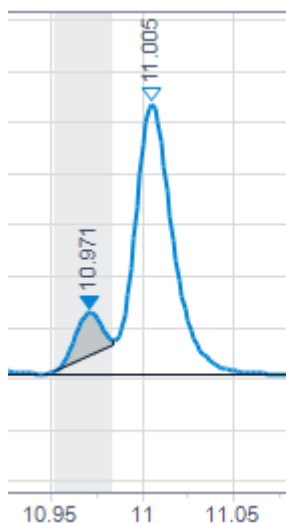
Area reject is ignored during manual integration.

Area% reject Sets the area% of the smallest peak of interest.

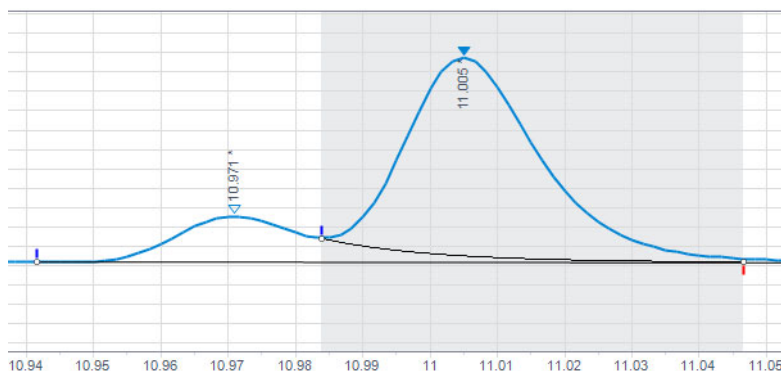
Any peaks with an area% less than the minimum area% are not reported. The integrator rejects any peaks with an area% smaller than the given value after baseline correction.

Enter the area% of the smallest peak expected. You can obtain this information by first integrating the data file with area and height reject set to zero (0). Use the **Area%** column in the integration results to choose an appropriate minimum value.

If a peak that is not integrated due to low area% is a rider peak, it will be merged with the parent peak.



If the parent peak is below the area% threshold, but the rider peak is above the threshold, the parent peak is kept, as the rider peak's calculation and baseline construction would otherwise be based on an excluded peak.



Height reject

Sets the height of the smallest peak of interest.

Any peaks that have heights less than this minimum height are not reported: The integrator rejects any peaks that are smaller than the **Height Reject** value after baseline correction.

NOTE

Height reject is ignored during manual integration.

Shoulders mode

Sets the initial method of detecting shoulders on peaks.

You can choose from:

Off	Shoulders are not detected.
Drop Baseline	Shoulders are integrated with a drop line.
Tangent Baseline	Shoulders are integrated with a tangent baseline.

This setting defines how the application handles peaks that are not baseline-separated. For more information on tangent skimming, see [“Tangent Skimming”](#) on page 38. If you use a tangent baseline, you can choose between different modes (see [“Tangent Skim Modes”](#) on page 42).

Choosing Peak Width

Choose the setting that provides just enough filtering to prevent noise being interpreted as peaks without distorting the information in the signal.

- To choose a suitable initial peak width for a single peak of interest, use the peak's time width at the base as a reference.
- To choose a suitable initial peak width when there are multiple peaks of interest, set the initial peak width to a value equal to or less than the narrowest peak width to obtain optimal peak selectivity.

Height Reject and Peak Width

Both **peak width** and **height reject** are very important in the integration process. You can achieve different results by changing these values.

- Increase both the height reject and peak width where relatively dominant components must be detected and quantified in a high-noise environment. An increased peak width improves the filtering of noise and an increased height reject ensures that random noise is ignored.
- Decrease height reject and peak width to detect and quantify trace components, those whose heights approach that of the noise itself. Decreasing peak width decreases signal filtering, while decreasing height reject ensures that small peaks are not rejected because they have insufficient height.
- When an analysis contains peaks with varying peak widths, set peak width for the narrower peaks and reduce height reject to ensure that the broad peaks are not ignored because of their reduced height.

Tuning Integration

It is often useful to change the values for the slope sensitivity, peak width, height reject, and area reject to customize integration. The figure below shows how these parameters affect the integration of five peaks in a signal.

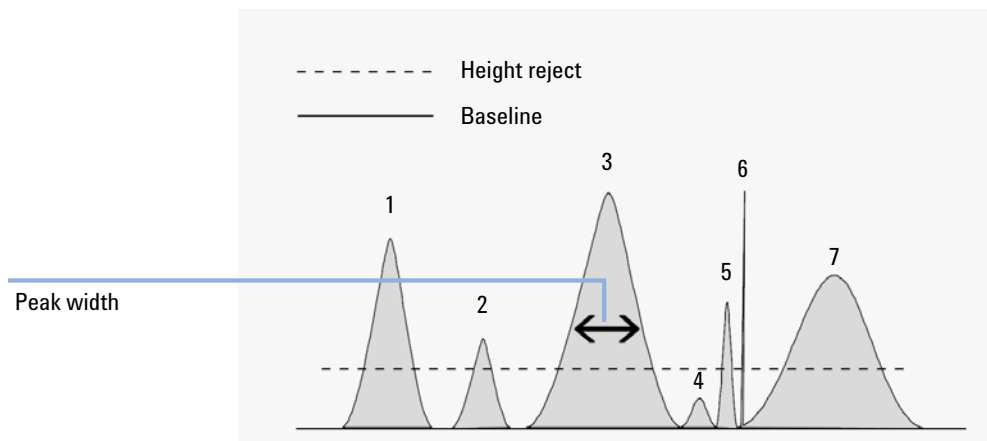


Figure 24 Using Initial Events

A peak is integrated only when all of the four integration parameters are satisfied. Using the peak width for peak 3, the area reject and slope sensitivity shown, only peaks 1, 3, and 7 are integrated.

- Peak 1** is integrated as all four integration parameters are satisfied.
- Peak 2** is rejected because the area is below the set area reject value.
- Peak 3** is integrated as all four integration parameters are satisfied.
- Peak 4** is not integrated because the peak height is below the Height Reject.
- Peak 5** is rejected because the area is below the set area reject value.
- Peak 6** is not integrated; filtering and bunching make the peak invisible.
- Peak 7** is integrated.

Table 5 Height and Area Reject Values

Integration Parameter	Peak 1	Peak 2	Peak 3	Peak 4	Peak 5	Peak 7
Height reject	Above	Above	Above	Below	Above	Above
Area reject	Above	Below	Above	Below	Below	Above
Peak integrated	Yes	No	Yes	No	No	Yes

Standard Integration Events: Timed Events

OpenLab CDS offers a set of timed events that allow a choice between the integrator modes of internal algorithm baseline definition and the user's definition. These timed events can be used to customize signal baseline construction when default construction is not appropriate. E.g. the user can create a new area sum event type (see **Area sum slice**), which does not alter the results of the default AreaSum. These events can be useful for summing final peak areas and for correcting short- and long-term baseline aberrations.

You can set the values either specifically for a given signal or globally for all signals.

Area reject See Initial Events ("[Standard Integration Events: Initial Events](#)" on page 46).

Area sum Sets points (**On/Off**) between which the integrator sums the areas.

The retention/migration time of a peak created with area summing is the average of the start and end times. If an **Area sum on** event occurs after the beginning of a peak but before the apex, the entire peak is included in the sum. If it occurs after the peak apex, but before the end of the peak, the peak is truncated and the area sum begins immediately.

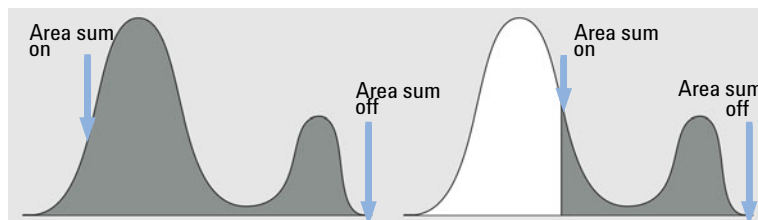


Figure 25 Area sum on event after peak apex, but before end of peak

If an **Area sum off** event occurs after the beginning of a peak but before the apex, the area sum ends immediately. The point on the signal where this occurs becomes a Valley Point. If the **Area sum off** event occurs after the apex, the event is postponed until the end of the peak.

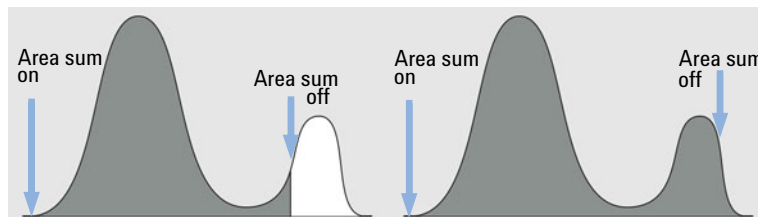


Figure 26 Area sum off event after beginning of peak, but before apex

Area sum slice This event allows you to define consecutive area sum intervals without any loss in area or time intervals.

This event is similar to the **Area Sum** event. However, with this event you can define contiguous area sum intervals without any loss in time intervals and integrated peak areas. A peak is split at the point where you set this event; area summing starts and ends exactly where the **Area Sum Slice** intervals are specified.

The retention time of the area sum slice peak is the middle of the slice time interval. The retention time does not change with identification or recalibration. It may only be shifted slightly, as the integrator only starts taking data points with the area sum slice start event, and ends with the area sum slice end event. Thus, the retention time may at most vary by the time between two data points.

Use the **Start** parameter to define the starting times for each area sum slice. The next start time is used as the end time for the preceding time-slice, so you can use several start events after each other.

The **Start-negA.** parameter defines the start of integration of a time-slice where any negative area (below the set baseline) is subtracted from the area of the time-slice.

The **End** parameter defines the end of the last time-slice. The area of the time-slice is calculated ignoring any area below the set baseline. If no other area sum slice events follow, the integrator resumes its regular peak detection again.

Within the range from a **Start** event to the next **End** event, the baseline is always one straight line with no changes in direction in between. Only after the end point (at least 0.001 min later) long term baseline changes can be applied again by using the events **Set Baseline from Range**, **Set Low Baseline from Range** or **Use Baseline from Range**.

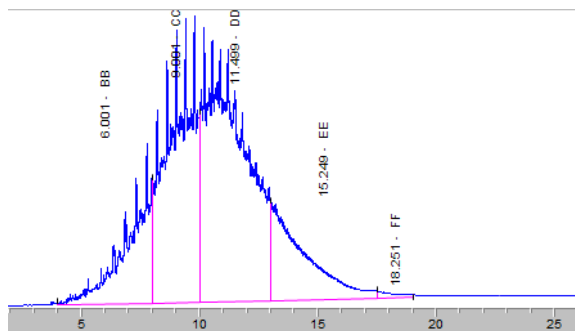


Figure 27 Example: Area Sum Slice

The figure above shows an example with the following timed events:

Table 6 Baseline construction

Time	Event	Parameter
4 min	Set BL from Range	+2 min
22 min	Set BL from Range	+4 min

Table 7 Area sum slices

Time	Event	Parameter
4 min	Area Sum Slice	Start
8 min	Area Sum Slice	Start
10 min	Area Sum Slice	Start
13 min	Area Sum Slice	Start
17.5 min	Area Sum Slice	Start
19 min	Area Sum Slice	End

Auto peak width

Turns on the automatic update of the peak width for the next peaks. It will resume with whatever the peak width is at that time and resume peak width tracking based on the previous found peak widths.

Baseline at valleys

Sets points (**On/Off**) between which the integrator resets the baseline at every valley between peaks.

The repeated resetting of the baseline can cut off corners of peaks. Such corners become negative area, they reduce the total measured area of the peaks.

This function is useful when peaks are riding on the back of a broad, low peak and you want the baseline to be reset to all the valley points.

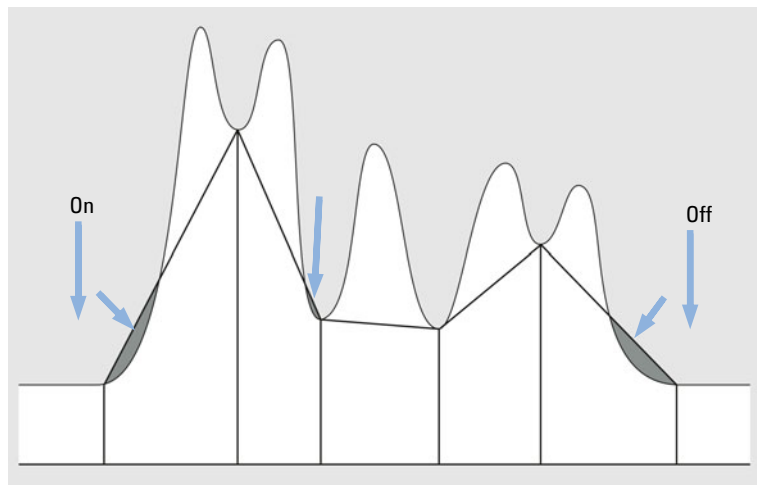


Figure 28 Baseline at valleys event

Baseline backwards

Sets a point at which the standard integrator extends the baseline, horizontally backward from the declared baseline point to this point.

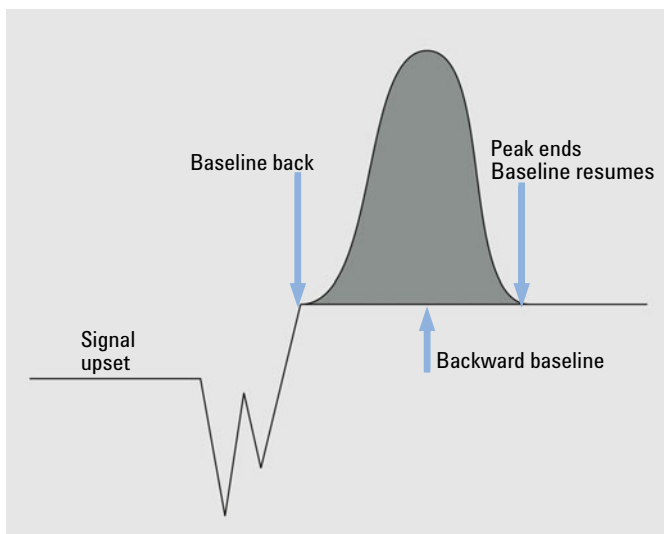


Figure 29 Baseline backwards event

Baseline hold

A horizontal baseline is drawn at the height of the established baseline from where the baseline hold event is switched on until where the baseline hold event is switched off.

Baseline next valley

Sets a point at which the integrator resets the baseline at the next valley between peaks, and then cancels this function automatically.

This function is useful in groups of merged peaks, which you assume are riding on the back or are in separate clusters close together. The function is ignored during area summing.

Baseline now

Sets a point (time) at which the integrator resets the baseline to the current height of the data point, if the signal is on a peak.

If the signal is on the baseline, the function is ignored and the detected baseline is used.

Detect shoulders

Sets points (**On/Off**) between which the integrator starts and stops detecting shoulders.

Shoulders are detected according to the specified **Shoulders Mode**. See "Standard Integration Events: Initial Events" on page 46.

- Fixed peak width** Sets the peak width and disables the automatic update of the peak width for the next peaks. To obtain good performance, set the peak width close to the width at half-height of the actual peaks.
- Height reject** See Initial Events ([“Standard Integration Events: Initial Events”](#) on page 46).
- Integration** Sets points (**On/Off**) between which the integrator stops and starts integrating. Peaks between the times where the integrator is turned off and on are ignored.

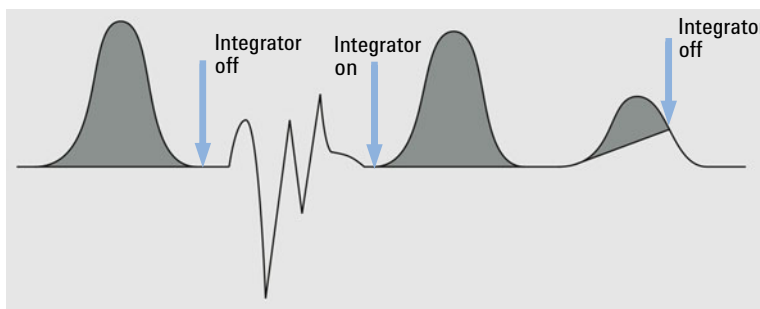


Figure 30 Integration event

The baseline is drawn from the last declared point including any resets for penetration. All other integrator functions together with set changes of peak width, threshold and area reject are ignored when the integrator is turned off. At the **On** and **Off** points, the baseline point is re-established.

When the integrator is set to restart, a new baseline point is reset at the current signal level.

This function is useful for ignoring parts of the chromatogram/electropherogram or to eliminate baseline disturbances.

- Maximum area** Sets the area of the largest peak of interest.

Any peaks that have areas greater than the maximum area are not reported: The integrator rejects any peaks that are greater than the maximum area value after baseline correction.

You can use this event, for example, to exclude the solvent peak of a GC chromatogram from the integration results, but include its rider peaks.

Maximum height

Sets the height of the largest peak of interest.

Any peaks that have heights greater than the maximum height are not reported: The integrator rejects any peaks that are higher than the maximum height value after baseline correction.

You can use this event, for example, to exclude the solvent peak of a GC chromatogram from the integration results, but include its rider peaks.

Negative peak

Sets points (**On/Off**) between which the integrator recognizes negative peaks.

When negative peaks are recognized, the integrator no longer automatically resets the baseline after penetration. From now on any penetration of the baseline will be integrated using the established baseline as zero. Areas are constructed relative to this baseline and are given an absolute value.

The negative peak function can only be used with confidence when the baseline drift is small compared with peak size, since the baseline is constructed from the declared baseline point at the start of the peak cluster up to the established baseline at the end of the peak.

NOTE

Area Summation is automatically deactivated if the **Negative Peaks On** event is activated.

Tangent skimming is also deactivated during negative peak detection; such peaks are separated by dropline.

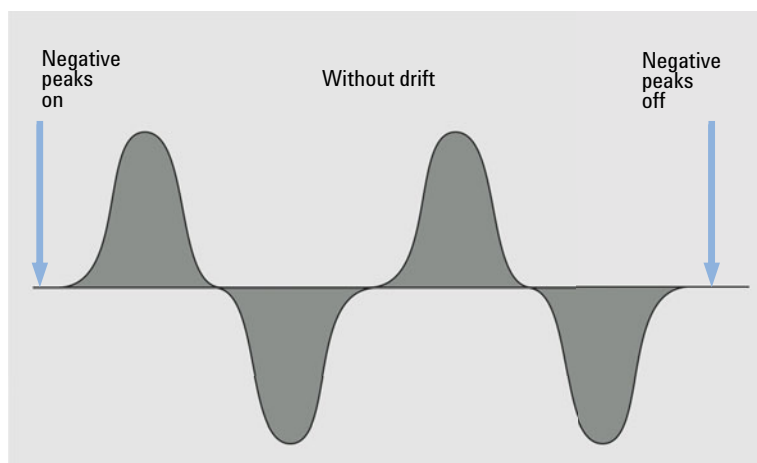


Figure 31 Negative peak event

Set baseline from range Uses a range of data points to calculate a statistically meaningful baseline point at the midpoint of a time-range.

The value that you provide is the time interval around a specified point in time. It defines the range to be used to determine the baseline point. See [“Baseline Correction Modes”](#) on page 35 for details of the statistical calculations of the baseline.

If you set the value =0, the nearest chromatogram data point is used as a baseline point; no statistics is done at all. If you set a negative value, the setting does the same as **Use baseline from range=Clear**. It stops the usage of the statistical baseline algorithm.

You can specify any area in the chromatogram for the baseline calculation. Ideally, it should be an area that is free from chemical background and contains only noise.

If you specify two **Set Baseline from Range** points (for example at the beginning and end of a chromatogram), the baseline between them is connected with a straight line.

Set low baseline from range Similar to **Set Baseline from Range**, but uses the lowest likely baseline point which allows 30 % more noise data points to be above it. Thus, baseline penetration is minimized.

Use **Set Low Baseline from Range** instead of **Set Baseline from Range** when the area of the chromatogram used for the calculation contains excessive chemical noise or electronic noise spikes.

Set Low Baseline from Range is calculated by a subtraction of one sigma (noise standard deviation) from the **Set Baseline from Range** y-value.

Shoulders mode See Initial Events ([“Standard Integration Events: Initial Events”](#) on page 46).

Slope sensitivity See Initial Events ([“Standard Integration Events: Initial Events”](#) on page 46).

Solvent peak Peaks above a specific slope in units of mV/s are detected as solvent peaks that lie outside the range of the analog-to-digital conversion.

The trailing peaks are automatically tangent-skimmed; you do not need to switch on the tangent skim event.

If solvent peak detection is off, droplines are drawn from the trailing peak instead of tangents.

Split peak Specifies a point at which to split a peak with a dropline.

NOTE

You cannot use **Split Peak** while **Area Sum** is switched on. To split a peak while **Area Sum** is switched on, use the corresponding manual integration event.

You cannot split skimmed peaks using the **Split Peak** event.

Tail tangent Specifies where to start or end tangent skimming.

skim

On

Sets a point at which the integrator sets a tangent skim on the trailing edge of the next peak. All peaks above the tangent are integrated to the reset baseline. The tangent is drawn from the valley before the small peak to the point after it where the detector signal gradient is equal to the tangent gradient. The tangent skim event time can be entered any time during the peak. Designates peak also as a solvent peak.

Off

Ends tangent skimming after current peak is completed or if in the designated interval no peaks are found (and a solvent will not inadvertently be designated in the next cluster).

Tangent skim
mode

The following models are available to calculate suitable peak areas:

- Exponential(Figure 20 on page 42)
- New exponential(Figure 21 on page 43)
- Standard(Figure on page 44)
- Straight(Figure 22 on page 44)

Unassigned peaks

With some baseline constructions, there are small areas that are above the baseline and below the signal, but are not part of any recognized peaks. Normally, such areas are neither measured nor reported. If unassigned peaks is turned on, these areas are measured and reported as unassigned peaks. The retention/migration time for such an area is the midpoint between the start and end of the area.

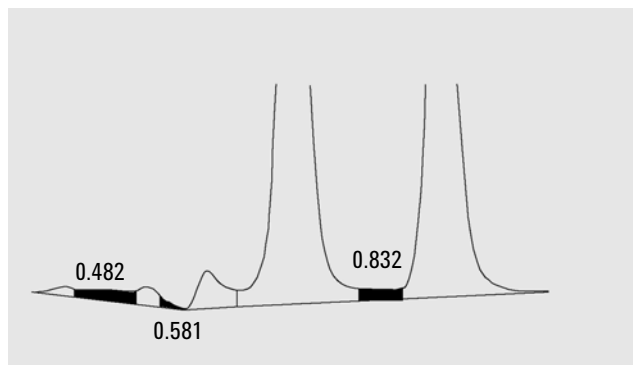


Figure 32 Unassigned Peaks

Update peak height

This event forces the integrator to use the absolute height of the highest data point as the peak height. Without this event, the maximum of an interpolated curve is used. The **Update peak height** event is useful especially in signals with extremely steep and angular peaks, or peaks preceded with a downward dip. Peaks like this are typical for MSD signals.

The start time of the **Update peak height** is not evaluated. The event always affects the entire chromatogram.

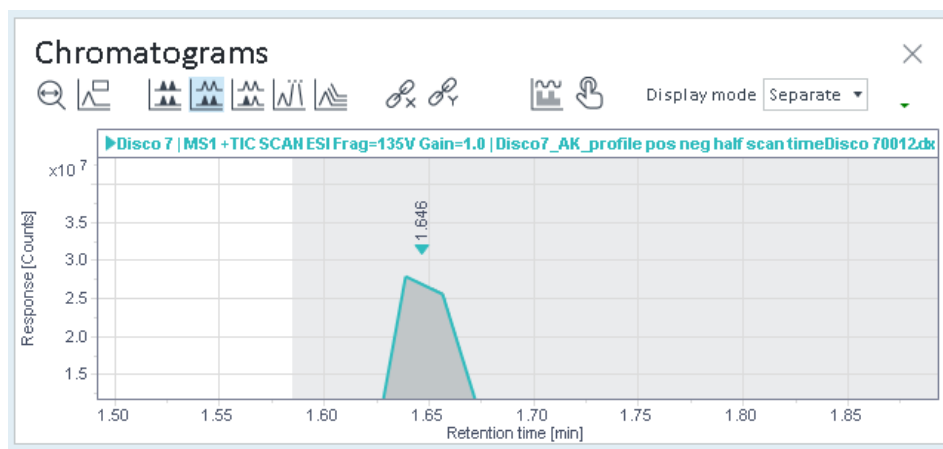


Figure 33 Default integration

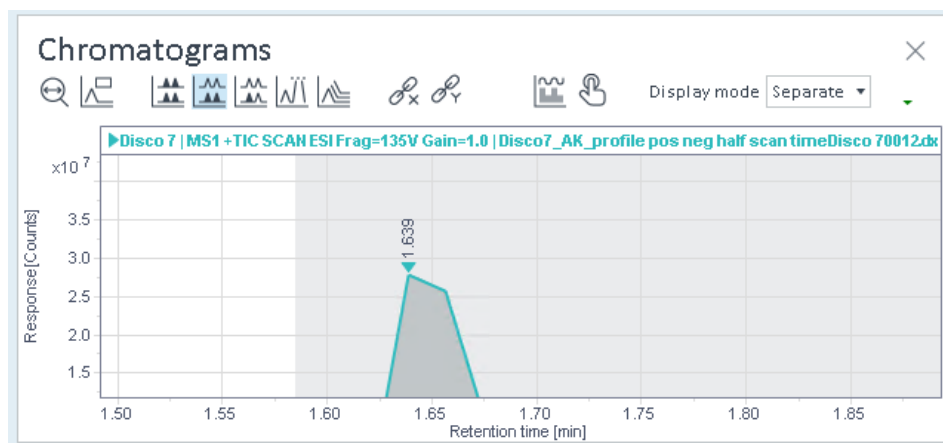


Figure 34 Integration with Update peak height event

Use baseline from range

Allows to project a baseline value to a later or earlier time to minimize baseline penetrations.

If the **Set Baseline from Range** or **Set Low Baseline from Range** value is calculated in an area with no chromatographic peaks, it can be advantageous to project the calculated baseline to the time immediately before the first peak of interest elutes (or to the time immediately after the last peak of interest has eluted). **Use Baseline from Range** allows you to make up to three such projections in either direction.

This event can be advantageous to use when you have constructed an upslope or downslope baseline, since otherwise the straight baseline might cut through the chromatogram curve unintentionally. The parameter tells the integrator from which of the baseline ranges to pick the baseline point and project the baseline to the baseline point at the given time interval.

You can use the following parameters:

- **Clear:** Clear the new baseline behavior and return to the traditional algorithm from this point.
- **Left:** Use the baseline value from the baseline range nearest to the left of this point in time.
- **Right:** Use the baseline value from the baseline range nearest to the right of this point in time.
- **Range 1—Range 9:** Use the baseline value from the given baseline range. Baseline ranges are counted from the beginning of the chromatogram.

See also the example under **Area Sum Slice** (Figure 27 on page 52).

Advanced Integration Events

The advanced integration events are provided for all signals.

Tangent skim mode

Define the type of baseline construction for peaks found on the upslope or downslope of a peak. See ["Tangent Skim Modes"](#) on page 42.

Exponential	Draws an exponential curve through the height-corrected start and end of each child peak.
New Exponential	Draws an exponential curve to approximate the trailing edge of the parent peak.
Standard	Combines exponential and straight line calculations for best fit.
Straight	Draws a straight line through the height-corrected start and end of each child peak.

Tail skim height ratio

Together with the **Skim valley ratio**, sets the conditions for tangent skimming a small peak on the tail of a solvent or other large peak. See ["Skim Criteria"](#) on page 39.

It is the ratio of the height of the baseline-corrected parent peak (H_p) to the height of the baseline-corrected child peak (H_c). Ratios higher than the specified value will enable skimming.

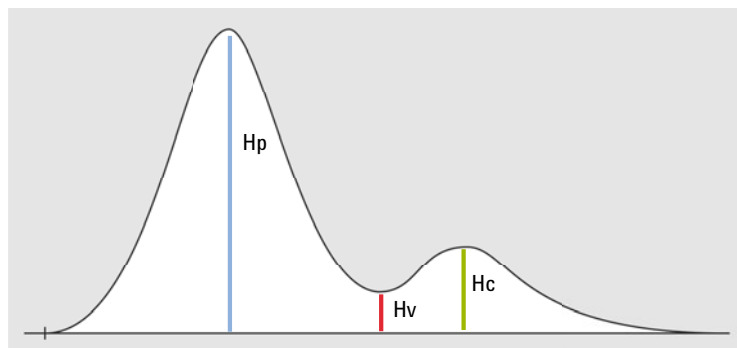


Figure 35 Example: Peak with tail skimming

Front skim height ratio

Together with the **Skim Valley Ratio**, sets the conditions for tangent skimming a small peak on the front of a solvent or other large peak. See “[Skim Criteria](#)” on page 39.

It is the ratio of the height of the baseline-corrected parent peak (H_p) to the height of the baseline-corrected child peak (H_c). Ratios higher than the specified value will enable skimming.

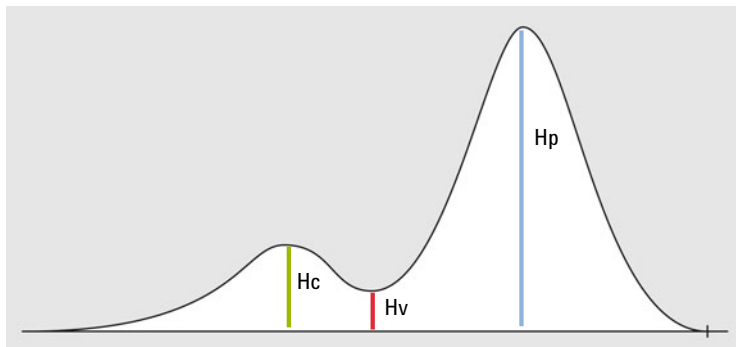


Figure 36 Example: Peak with front skimming

Skim valley ratio

Together with the **Tail Skim Height Ratio** or **Front Skim Height Ratio**, sets the conditions for tangent skimming a small peak on the tail or front of a solvent or other large peak. See “[Skim Criteria](#)” on page 39.

It is the ratio of the height of the baseline-corrected child peak (H_c) to the height of the baseline-corrected valley (H_v). Ratios lower than the specified value will enable skimming.

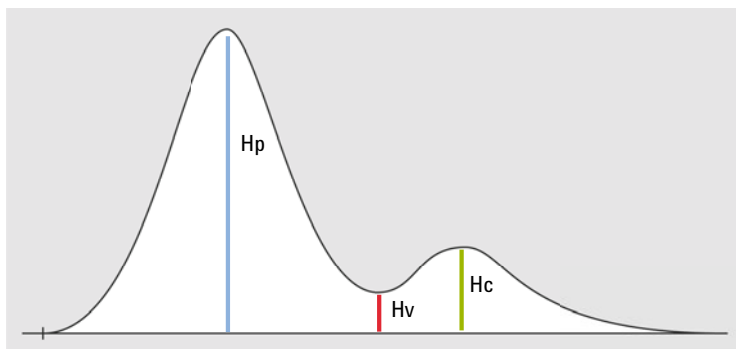


Figure 37 Example: Peak with tail skimming

**Baseline
correction
mode**

Sets the type of baseline correction. See “Baseline Correction Modes” on page 35.

You can choose between the following parameters:

Classical	Accepts baseline penetrations.
No penetrations	Removes baseline penetrations by reconstructing the baseline.
Advanced	The integrator tries to optimize the start and end locations of the peaks, re-establishes the baseline for a cluster of peaks and removes baseline penetrations.

**Peak-to-Valley
ratio**

Used to decide whether two peaks that do not show baseline separation are separated using a drop line or a valley baseline, it is the ratio of the baseline-corrected height of the smaller peak to the baseline-corrected height of the valley. See “Peak-to-Valley Ratio” on page 37.

When the peak to valley ratio is lower than the specified value, a drop line is used (A); otherwise, a baseline is drawn from the baseline at the start of the first peak to the valley, and from the valley to the baseline at the end of the second peak (B).

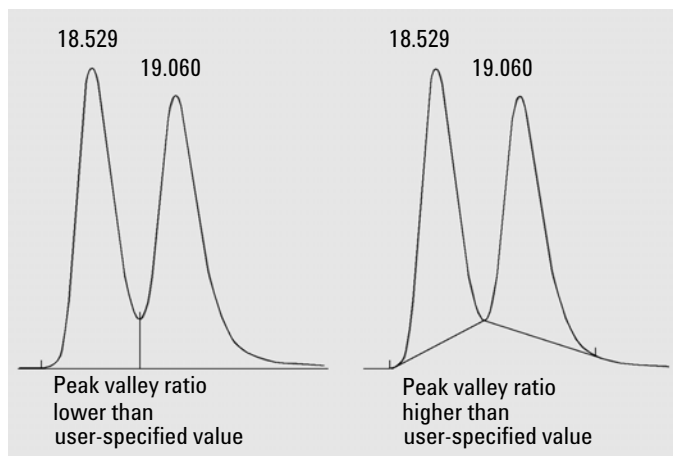


Figure 38 Effect of peak valley ratio on the baselines



3 Integration with EZChrom Integrator

Integration Events 67

Baseline Code Descriptions 83

This chapter contains the description of EZChrom integration events.

Integration Events

You can set the values either specifically for a given signal or globally for all signals. To add a timed event, right-click in the parameters table.

There are different types of integration events: For some of them, you can define a time range with start and stop time during which a parameter is active. For others, you can define a specific value to be used from a start time or during a time range. The columns **Time Stop [min]** and **Value** are enabled or grayed out, depending on the type of event.

NOTE

Other than in OpenLab EZChrom, the stop time is grayed out for the following integration events in OpenLab CDS:

- Width
- Threshold
- Shoulder sensitivity
- Minimum area
- Reset baseline
- Reset baseline at valley

Width Used to distinguish true peaks from noise. The system uses the default value of width = 0.2 min.

The **Width** event is used to calculate a value for bunching, or smoothing, the data points before the integration algorithm is applied. Integration works best when there are 20 points across a peak. If a peak is over sampled (i.e. the sampling frequency was too high), the **Width** parameter will be used to average the data such that the integration algorithm sees only 20 points across the peak.

A **Width** event will be applied to a given peak as long as it occurs before or on the apex of the peak.

The **Width** parameter is only used to correct for over-sampling. It cannot correct for data that was under-sampled (i.e. sampling frequency too low causing fewer than 20 points acquired across the narrowest peak).

The diagrams below show examples of how incorrect values can effect the peak baseline.

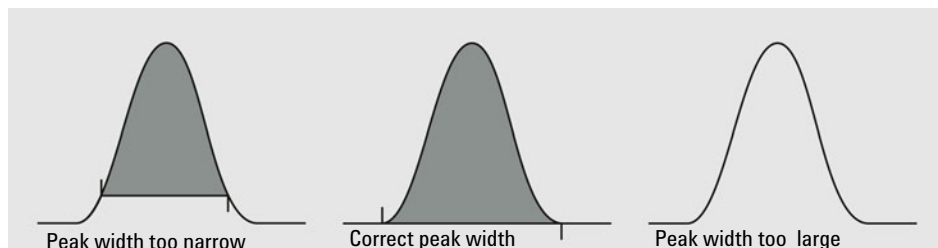


Figure 39 Width

NOTE

In most circumstances, an initial Width value based on the narrowest peak in the chromatogram will be adequate for proper integration of all peaks. However, a new Width timed event should be entered every time a peak width doubles.

NOTE

Extreme values of both Width and Threshold (too large or too small) can result in peaks not being detected.

Threshold

This parameter is the first derivative, used to allow the integration algorithm to distinguish the start and stop of peaks from baseline noise and drift. The **Threshold** value is based on the highest first derivative value determined in a section of the chromatogram.

The diagrams below show examples of how incorrect values can effect the peak baseline.

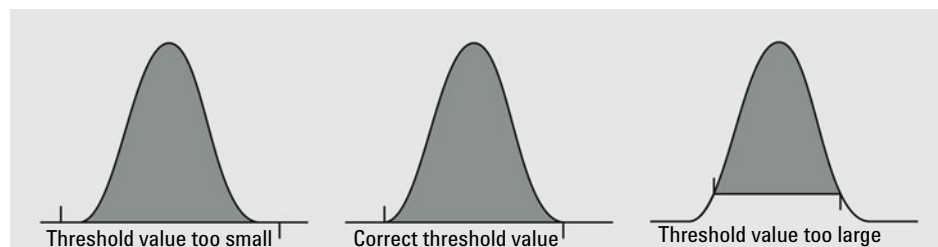


Figure 40 Threshold

NOTE

Extreme values of both Width and Threshold (too large or too small) can result in peaks not being detected.

Shoulder sensitivity

This parameter is used to enable the detection of shoulders on larger peaks. A larger value will decrease shoulder sensitivity while smaller values increase sensitivity to shoulder peaks. The **Shoulder Sensitivity** value is based on the highest second derivative value determined in a section of the chromatogram.

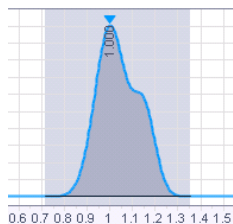


Figure 41 Shoulder sensitivity value set too high

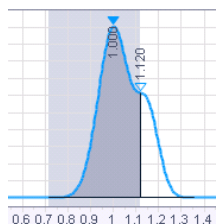


Figure 42 Shoulder sensitivity value set correctly

Integration off This event turns off the integration of your chromatogram during the range specified. This event is useful if you are not interested in certain areas of your chromatogram, and do not wish peaks to be reported for that section.

When using **Integration Off** to disable peaks, these regions will be included in the noise calculation. Leave all peaks integrated to get the correct noise values.

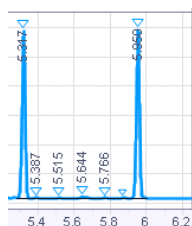


Figure 43 Default integration

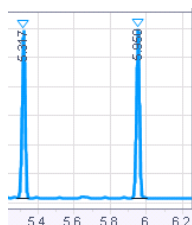


Figure 44 Integration off from 5.35 to 5.85 min

Valley to valley

This event causes the baselines of peaks that are not totally resolved (i.e. do not return to baseline) to be drawn to the minimum point between the peaks.

If this event is not used, a baseline is projected to the next point at which the chromatogram returns to baseline, and a perpendicular is dropped for peaks which do not reach baseline.

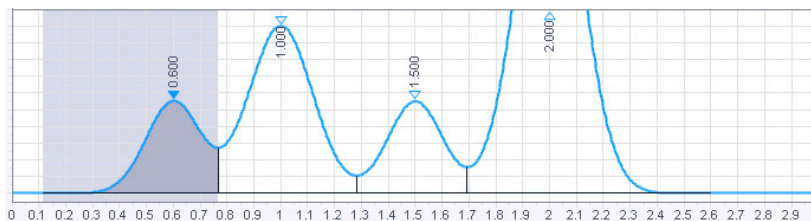


Figure 45 Default integration

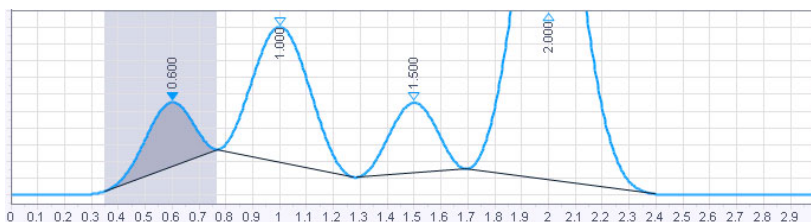


Figure 46 Integration with Valley to valley event

**Horizontal
baseline**

This event allows you to project the baseline forward horizontally for the peaks within the specified time range. The baseline starts where the first peak within the defined time range begins. The end of the baseline is where it intersects the signal, or where the last peak ends.

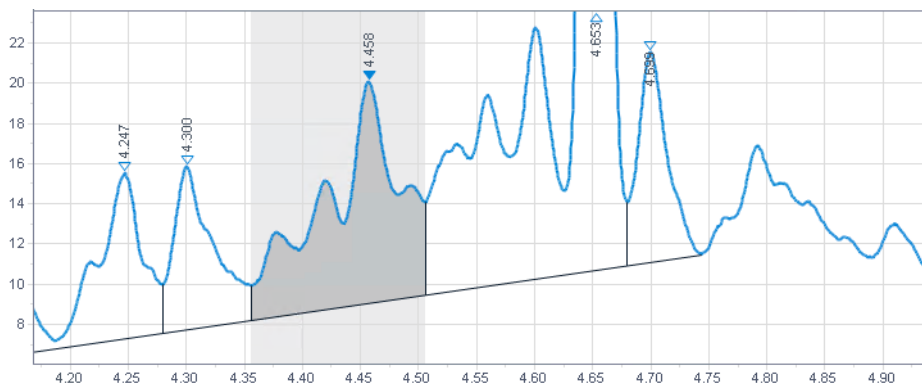


Figure 47 Default integration

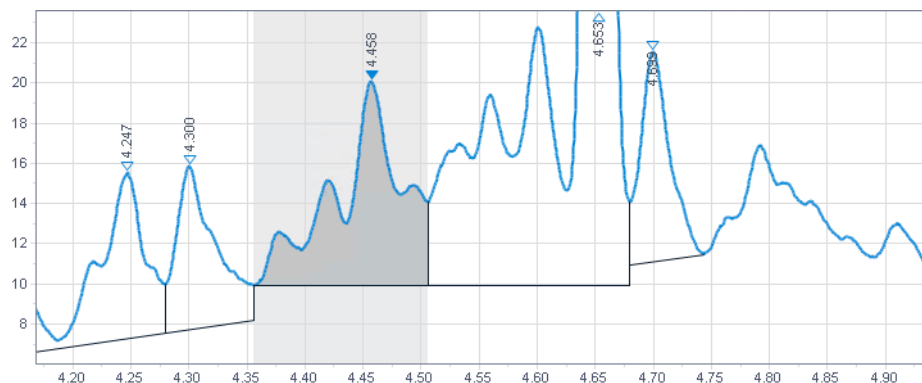


Figure 48 Integration with Horizontal baseline event between 4.4 and 4.69 min

Backward horizontal baseline

This event is used to force a horizontal baseline in the direction of the beginning of the chromatogram. A backward horizontal baseline will be created for the peaks within the specified time range. The baseline is defined by the end of the last peak within the defined time range ends. It is projected back to where it intersects the signal, or where the first peak begins.

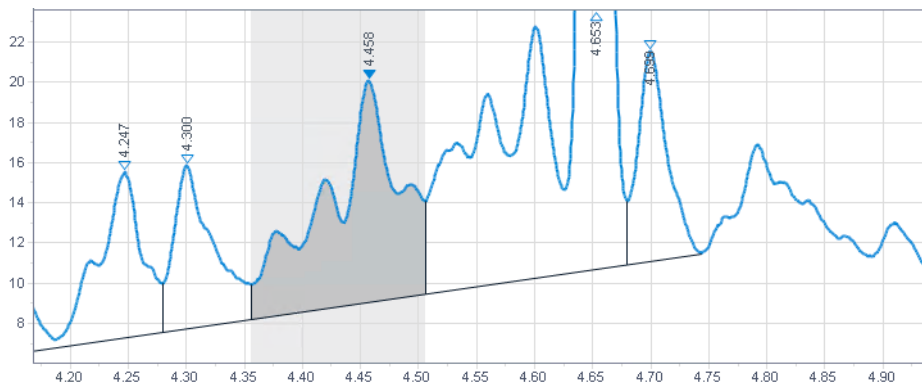


Figure 49 Default integration

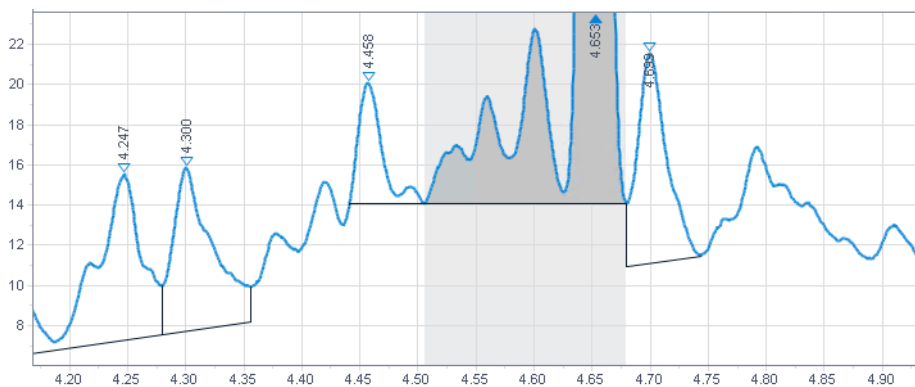


Figure 50 Integration with Backward horizontal baseline event between 4.4 and 4.69 min

Lowest point horizontal baseline

This event is applied to all peaks whose retention times fall within the time range defined by **Time (min)** and **Time Stop (min)**.

The **Value** defines the time where the integrator starts to search for the lowest baseline point.

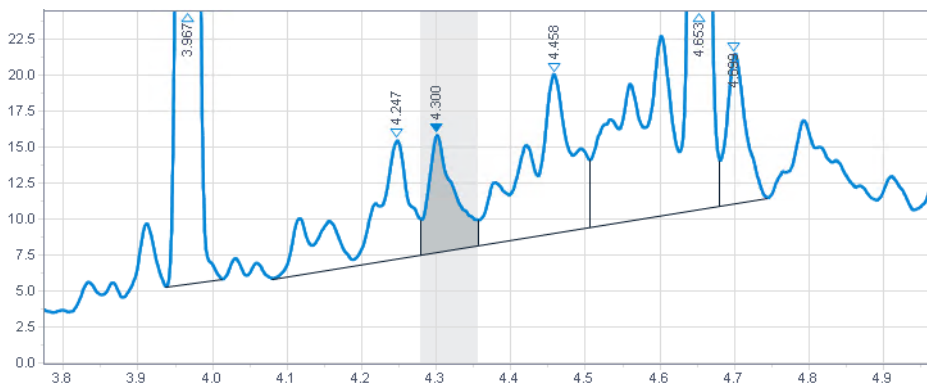


Figure 51 Default integration

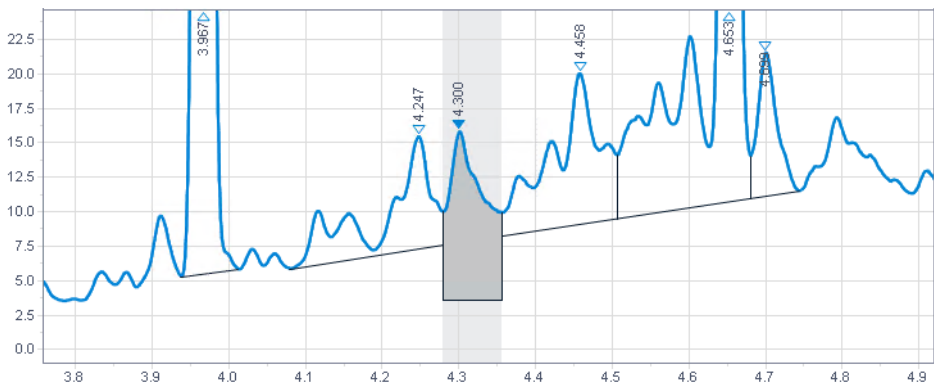


Figure 52 Integration with Lowest point horizontal baseline event between 4.25 and 4.35 min, with Value set to 3.8 min

Tangent skim This event is used to integrate a small peak located on the tailing edge of a larger peak. The baseline of the small peak becomes a tangent drawn from the valley of the larger peak to the tangent point on the chromatogram.

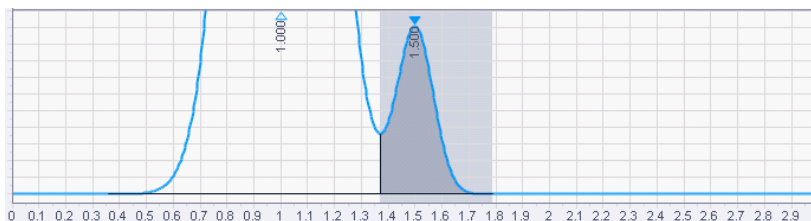


Figure 53 Default integration

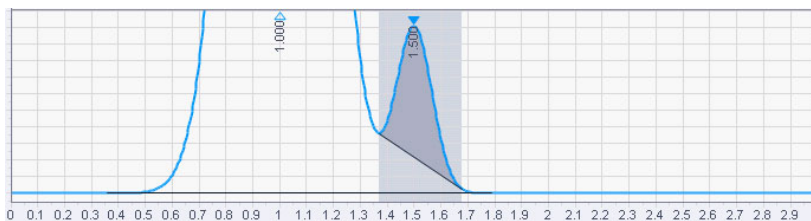


Figure 54 Integration with Tangent skim event

Front tangent skim This event is used to force a tangential baseline for a daughter peak on the leading edge of a mother peak.

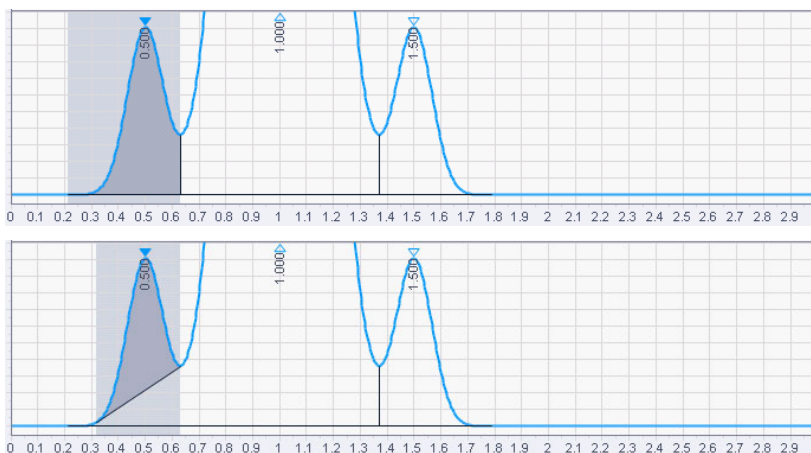


Figure 55 Integration with Front tangent skim event

Exponential skim

This event is used to integrate small peaks located on the tailing edge of a larger peak. The baseline of the small peak becomes an exponential drawn from the valley of the larger peak to the tangent point on the chromatogram.

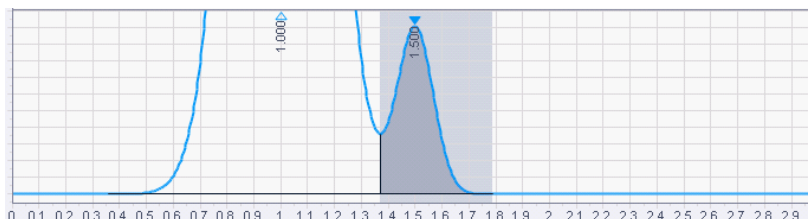


Figure 56 Default integration

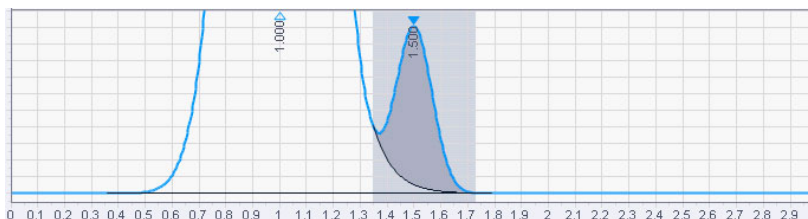


Figure 57 Integration with Exponential skim event

Front exponential skim

This event is used to force an exponential baseline for a daughter peak on the leading edge of a mother peak.

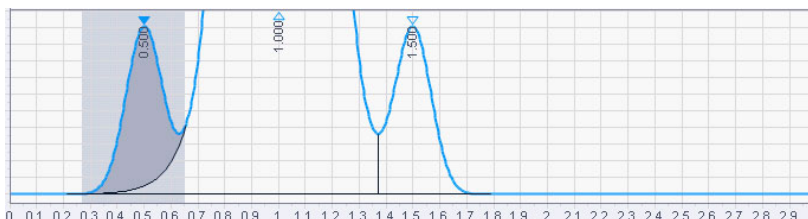
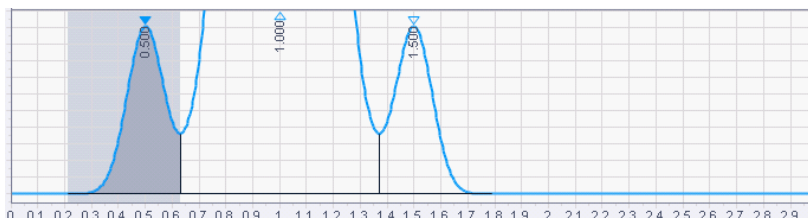


Figure 58 Integration with Front exponential skim event

Minimum area

This event allows you to enter an area limit for peak detection. Peaks whose areas fall below this minimum area will not be integrated and reported as peaks. This event is useful for eliminating noise or contaminant peaks from your report.

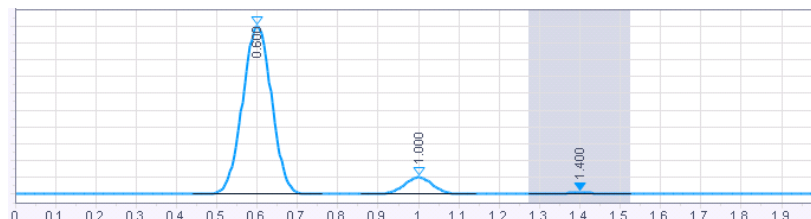


Figure 59 Default integration

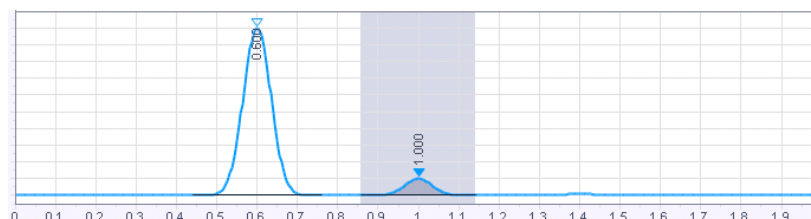


Figure 60 Integration with Minimum area event

Negative peak

This event causes portions of the chromatogram that drop below the baseline to be integrated using the normal peak logic and reported as true peaks. This event is useful when using detectors such as Refractive Index types which give a negative response to certain compounds.

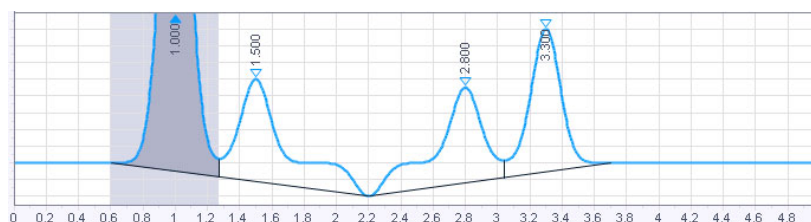


Figure 61 Default integration

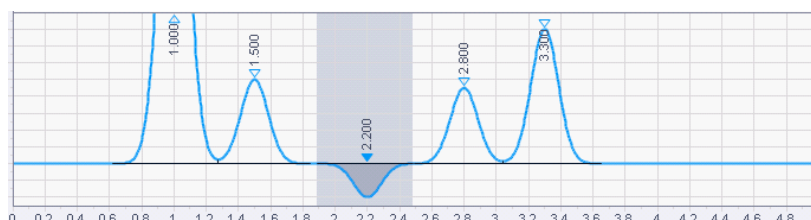
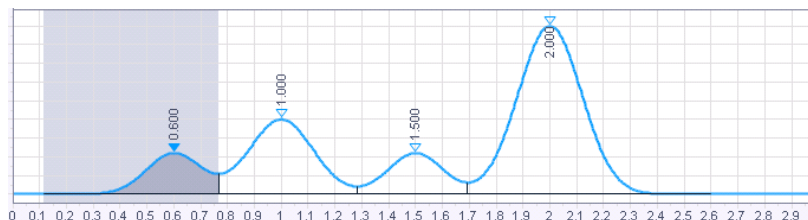
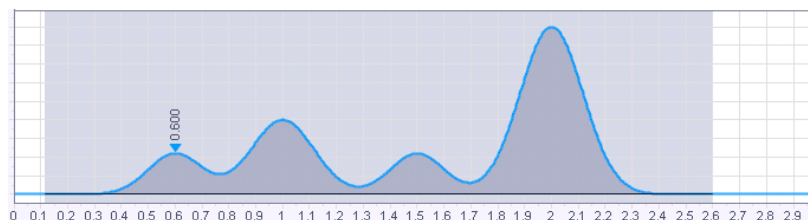


Figure 62 Integration with Negative peak event

Disable end of peak detection

This event is used to turn off end of peak detection between the specified times, forcing the software to treat peaks falling within the window of the event as a single peak. This event is a useful way to combine the areas of a series of contiguous peaks into one area. Because the peaks are considered to be part of a single peak, the retention time is assigned to the time of the first apex after the **Disable End of Peak Detection** event.

**Figure 63** Default integration**Figure 64** Disable end of peak detection between 0.4 and 2.3 min

Manual baseline

This event allows you to change the way the baseline for a peak is drawn without changing the integration parameters. The baseline will be drawn from the signal at the start time to the signal at the stop time.

This is convenient when you want to change where a baseline is drawn for a peak without changing how the baseline is drawn for other peaks in the chromatogram.

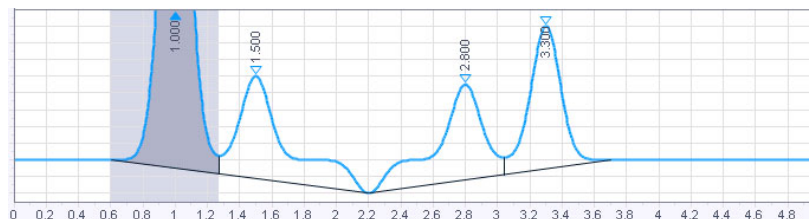


Figure 65 Default integration



Figure 66 Integration with manual baseline between 2.3 and 3.6 min

Manual peak This command allows you to define the start and stop time of a peak that was not previously detected. This is convenient when you want to force integration of a peak, but do not want to change your overall integration parameters.

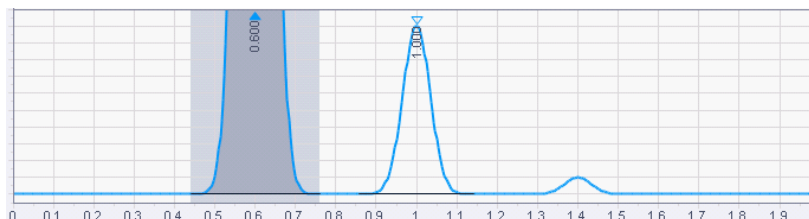


Figure 67 Default integration

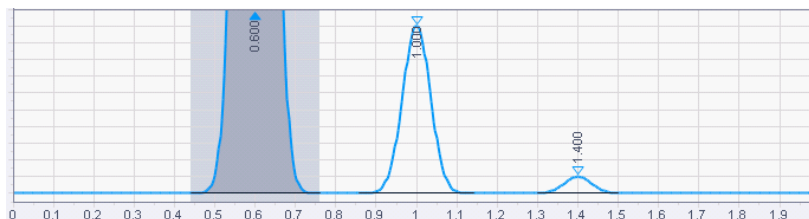


Figure 68 Small peak integration forced using Manual peak event between 1.3 and 1.5 min

Split peak This event is used to force a perpendicular drop-line integration in a peak. The perpendicular will be dropped at the time where the event is inserted.

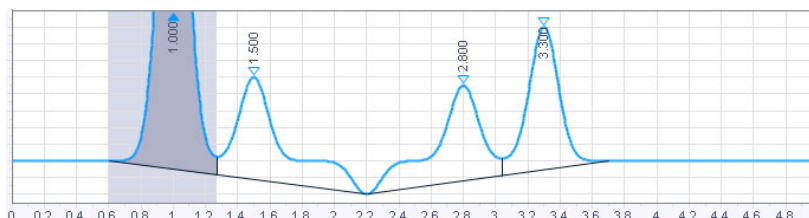


Figure 69 Default integration

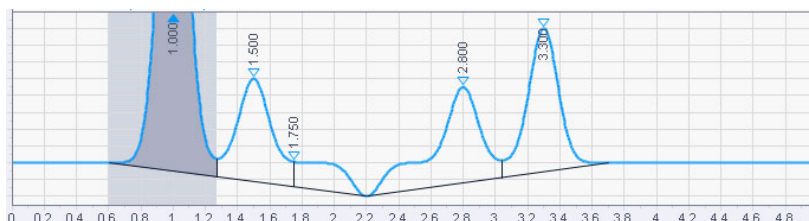


Figure 70 Integration with a Split peak event at 1.75 min

Force peak start / Force peak end

These events are used to force the start or stop of the peak integration to a specific point.

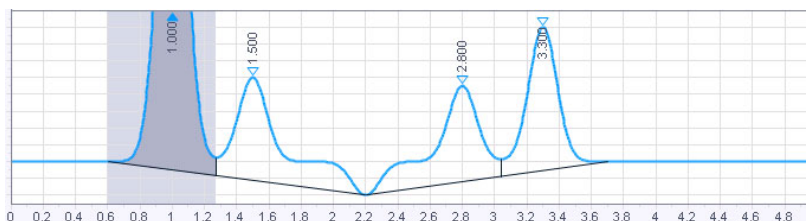


Figure 71 Default integration

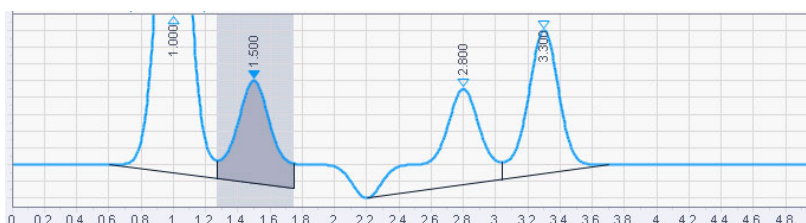


Figure 72 Integration with a Force peak end event at 1.75 min

Reset baseline

This event lets you set the baseline at a designated point on the chromatogram.

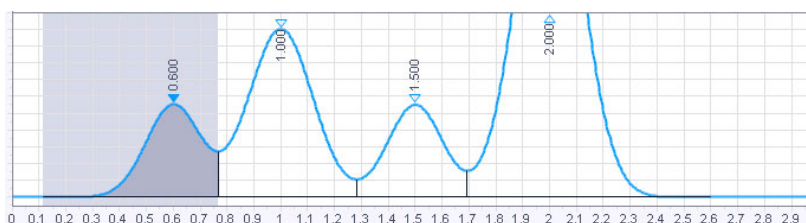


Figure 73 Default integration

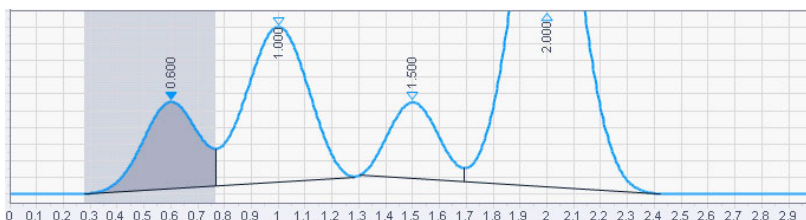


Figure 74 Integration with a Reset baseline event at 1.3 min

Reset baseline at valley

This event will cause the baseline to be reset at the next valley detected after the event.

NOTE

The event should be placed after the start of the first peak in the cluster; otherwise the start of the peak will be identified as the valley.

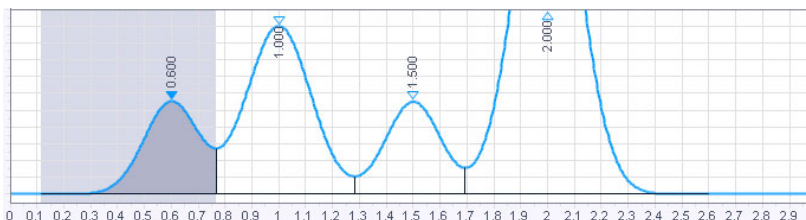


Figure 75 Default integration

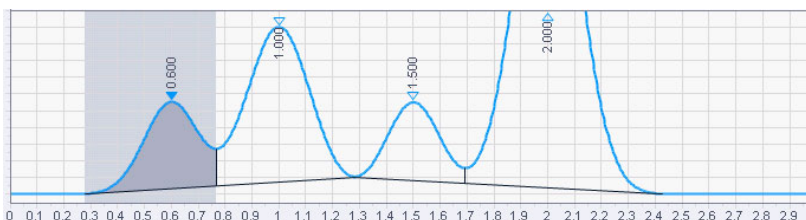


Figure 76 Integration with a Reset baseline at valley event at 1.2 min

Maximum area

Sets the area of the largest peak of interest.

Any peaks that have areas greater than the maximum area are not reported: The integrator rejects any peaks that are greater than the maximum area value after baseline correction.

You can use this event, for example, to exclude the solvent peak of a GC chromatogram from the integration results, but include its rider peaks.

Maximum height

Sets the height of the largest peak of interest.

Any peaks that have heights greater than the maximum height are not reported: The integrator rejects any peaks that are higher than the maximum height value after baseline correction.

You can use this event, for example, to exclude the solvent peak of a GC chromatogram from the integration results, but include its rider peaks.

Baseline Code Descriptions

Baseline Code Descriptions

A baseline code consists of two letters. The first letter denotes the peak beginning baseline type and the second letter indicates the peak ending baseline type. The baseline codes are included in the **Injection Results** table and in all default report templates.

Injection Results							
Peaks		Summary					
#	Name	BL code	Signal description	RT (min)	Area	Area%	Height
1		VB	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	1.940	356.576	27.285	65.775
2		BB	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	2.578	314.315	24.051	49.125
3		BB	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	3.022	333.322	25.506	59.034
4		BV	DAD1A,Sig=272.0,16.0 Ref=380.0,20.0	4.517	302.650	23.159	49.419

Figure 77 Injection results

Signal: DAD1A,Sig=272.0,16.0 Ref=380.0,20.0					
RT [min]	Type	Width [min]	Area	Height	Area%
1.940	VB	0.53	356.58	65.78	27.28
2.578	BB	0.43	314.31	49.12	24.05
3.022	BB	0.64	333.32	59.03	25.51
4.517	BV	0.59	302.65	49.42	23.16
		Sum	1306.86		

Figure 78 Example: Table from Short Area report

B	Baseline
C	Exponential
f	Force Peak Start or Stop (user defined)
I	Peak ended by Integration Off event
N	Begin negative peak
P	End negative peak
H	Forward horizontal
h	Backward horizontal
M	Manual baseline or Manual peak
m	Move baseline Start/Stop
S	Shoulder
T	Tangent skim
V	Valley
v	Forced valley point
x	Split peak
E	End of chromatogram encountered before the end of peak was found. End of chromatogram used as peak end.
R	Reset Baseline
L	Lowest Point Horizontal

4

Peak Identification

What is Peak Identification?	86
Evaluation of the Retention Time Window	87
Conflict Resolution	88
Relative Retention Times	89
Calculations for Relative Retention Times (RRT)	90
Time Reference Compound	91
About Time Reference Compounds	91
Calculations for Time Reference Compounds	92
Update Processing Method	93
Retention Time Updates	93
Calculations for Updated Retention Times	94
Example: Retention Time Updates with RRT	95
Calculation for Global Retention Time Shift	97

This chapter describes the concepts of peak identification.

What is Peak Identification?

Peak identification identifies the compounds in an unknown sample based on their chromatographic characteristics.

The identification of these compounds is a necessary step in quantitation if the analytical method requires quantitation. It is possible to create a valid processing method with identification even without quantitation. The signal characteristics of each component of interest are stored in the compound table of the method.

The function of the peak identification process is to compare each peak in the signal with the peaks stored in the compound table.

The identification is based on expected retention time, absolute retention time window, and relative retention time window in %. The final retention time window is the sum of relative and absolute windows, applied symmetrically to the expected retention time.

The expected retention time is either specified in the method as absolute time value or calculated from a relative retention time. Time reference compounds may be used to correct the expected retention times based on possible shifts observed by specific reference compounds.

$$\text{Wnd Wdth} = \text{Abs R T Wnd} + \frac{\text{Exp R T} * \text{Rel R T Wnd}}{100}$$

where

Abs R T Wnd	Absolute retention time window
Exp R T	Expected retention time
Rel R T Wnd	Relative retention time window
Wnd Wdth	Window width

$$\text{R T Wnd} = [\text{Exp R T} - \text{Wnd Wdth}; \text{Exp R T} + \text{Wnd Wdth}]$$

where

Exp R T	Expected retention time
R T Wnd	Retention time window
Wnd Wdth	Window width

Evaluation of the Retention Time Window

The identification window is the sum of relative and absolute window, applied symmetrically to the expected retention time. For example:

Expected retention time = 1 min

Absolute retention time window = 0.2 min

Relative retention time window = 10 % = 1 min * 10/100 = 0.1 min

Identification window = $[1 - 0.2 - 0.1 ; 1 + 0.2 + 0.1] = [0.7 ; 1.3]$

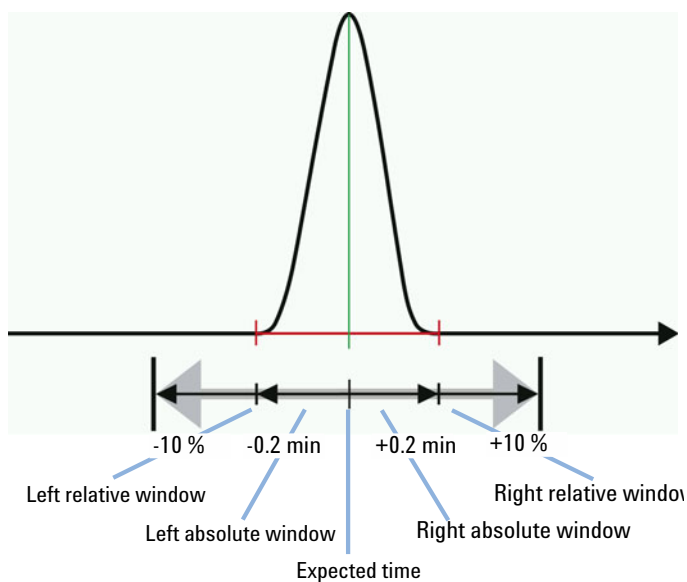


Figure 79 Identification window

Conflict Resolution

If multiple peaks are within the retention time window, there are different ways how to identify a particular peak. You can choose between the following values for **Peak match** in the compound identification parameters:

- **First:** Use the first peak in the retention time window.
- **Last:** Use the last peak in the retention time window.
- **Closest** (default setting): Use the peak that is closest to the expected retention time.
- **Largest area:** Use the peak with the largest area in the retention time window.
- **Largest height:** Use the peak with the largest height in the retention time window.

If the conflict cannot be resolved, none of the peaks will be identified and a warning will be written into the processing log.

Relative Retention Times

You can use relative retention times to check if the identification of your compounds is correct. The retention time of a compound is compared to the retention time of another specific given compound (also referred to as RRT reference). The ratio of the two retention times, that is, the relative retention time RRT, is normally a known number which you can provide in the application.

The RRT values themselves have no impact on the compound identification. Only the absolute expected retention times are used for this purpose. They are either specified in the processing method as absolute time values or calculated from relative retention times. Time reference compounds or method updates may be used to correct these absolute retention time windows based on possible shifts.

The following example shows the identification parameters for an RRT reference compound with an associated compound.

Compound Table			Qualifier Setup	General			
#	Type	Name	Signal	Is RRT ref.	Associated RRT reference	Exp. RT (min)	RRT
1		Ref #1	DAD1A	☒		2.000	1.000
2		Compound #1	DAD1A	☐	Ref #1 (DAD1A)	3.000	1.500

Associated compound

RRT reference compound

If you change the expected RT of the associated compound, its RRT value will automatically be recalculated. Also vice versa, if you change its RRT value, its expected RT will be recalculated.

If you change the expected RT of the RRT reference compound, the system recalculates the expected RT of the associated compounds.

If you use *time reference compounds* with RRT reference compounds, the retention time shift is applied to the RRT reference compound (see ["Calculations for Time Reference Compounds"](#) on page 92). The system recalculates the expected RT of the associated compound so that the RRT values do not change.

Compound Table		Qualifier Setup		General					
#	Type	Name	Signal	Is time ref.	Associated time ref.	Exp. RT (min)	RRT	Is RRT ref.	
1		Ref #1	DAD1A	☐		2.000	1.000	☒	
2		Compound #1	DAD1A	☐	Time Ref (DAD1A)	3.000	1.500	☐	
3		Time Ref	DAD1A	☒		3.500		☐	

RT: Adjusted, if required

RRT: Remains constant

You can also use the **Update RT** function with RRT reference compounds. However, you can only configure the update parameters for the RRT reference compounds. The associated compounds are forced to use the same values as their references, so that the RRT values do not change.

Compound Table		Qualifier Setup		General				
#	△	Type	Name	Signal	Is RRT ref.	Associated RRT reference	RT update	RT update factor (%)
1		RRT Ref	DAD1A	☒			After each run	50.000
2		Compound #1	DAD1A	☐		RRT Ref (DAD1A)	After each run	

Associated compound uses same RT Update values as RRT reference

RRT Reference compound uses RT update function

For *not identified peaks*, there is no associated compound. In this case, the first RRT reference compound is used to calculate the RRT value.

Calculations for Relative Retention Times (RRT)

Calculation of relative retention time (RRT) from expected retention times:

$$\text{RRT} = \text{expected RT}_{\text{compound}} / \text{expected RT}_{\text{reference}}$$

If you change the RRT value of the associated compound, its expected RT is recalculated as follows:

$$\text{Expected RT}_{\text{compound}} = \text{RRT} * \text{expected RT}_{\text{reference}}$$

Time Reference Compound

About Time Reference Compounds

If you use time reference compounds, the application corrects the absolute retention time windows based on possible shifts observed by specific reference compounds.

One or more compounds in the processing method can be marked as time reference compounds. For each compound or timed group, a time reference compound can be selected to correct the expected retention time. The extent of correction can be adjusted by an individual correction factor, which can be selected for each compound for correcting the expected retention time (column **Factor**, default = 1).

If a compound has a time reference compound assigned to it, the expected retention time will be corrected by the shift of the assigned time reference compound. The compound identification algorithm will use the corrected expected retention time for identifying the peak in the chromatogram. In case of timed groups, the time ranges are corrected by the shift. Generally the shift is corrected by the entered correction factor. If an associated reference compound is not found in the chromatogram, the linked peaks and time groups are not identified.

If *internal standards* are used and **Use time reference compounds** is selected, the internal standards are by default set as time reference compounds.

If you use *time reference compounds* and RRT reference compounds, both expected $RT_{\text{reference}}$ and expected RT_{compound} are adjusted, so that the RRT remains constant.

NOTE

When using time reference compounds, *all* compounds and timed groups must have a time reference compound assigned to them. Otherwise the method is inconsistent and cannot be used for reprocessing.

Calculations for Time Reference Compounds

If you use time reference compounds, the application corrects the absolute retention time windows based on possible shifts observed by selected time reference compounds.

Shift of the time reference compound

$$\text{Shift}_{\text{Ref}} = \text{ActualRT}_{\text{Ref}} - \text{ExpRT}_{\text{Ref}}$$

where

$\text{Shift}_{\text{Ref}}$ Time shift of the reference compound

$\text{ActualRT}_{\text{Ref}}$ Actual retention time of the reference compound

$\text{ExpRT}_{\text{Ref}}$ Expected retention time of the reference compound

RT of associated compound

For compounds that use a time reference, the expected retention time is calculated using an additional factor.

$$\text{CorrectedExpRT} = \text{ExpRT} + (\text{Shift}_{\text{Ref}} * \text{Factor})$$

where

CorrectedExpRT Corrected expected retention

ExpRT Expected retention time

$\text{Shift}_{\text{Ref}}$ Time shift of time reference compound

Factor Factor for compounds with associated time reference compound (**Factor**)

**Start and stop
time of
associated
timed group**

For timed groups, the expected start and stop times are calculated accordingly:

$$\text{Corrected Range Start} = \text{Range Start} + (\text{Shift}_{\text{Ref}} * \text{Factor})$$

$$\text{Corrected Range End} = \text{Range End} + (\text{Shift}_{\text{Ref}} * \text{Factor})$$

where

Corrected Range Start Corrected start time of timed group

Range Start Start time of timed group

Corrected Range End Corrected stop time of timed group

RangeEnd Stop time of timed group

Shift_{Ref} Time shift of time reference compound

Factor Factor for timed groups with an associated time reference compound (**Factor**)

Update Processing Method

Retention Time Updates

Based on the retention time update type (**Never**, **After each run**, or **After calibration standards**) of all identified compounds or timed groups, the expected retention time in the processing method is automatically updated after the corrected expected retention time has been calculated. If the compound can be found based on the corrected value, the corrected value becomes the new expected value in the method.

Retention time updates can be applied with or without time references.

NOTE

If retention time update is set to **After each run** or **After calibration standards** all injections are processed in sequential order. The change in the method will be applied with the next injection and no more parallel processing of non-calibration injections can be done.

In addition to updating the retention times during the run, you can also manually shift all retention times by a given value.

Calculations for Updated Retention Times

To correct the expected retention times, the application reads the current retention time and calculates the shift to the expected retention time. This shift, multiplied by a compound-specific weighting factor, is added to the expected retention time.

$$\text{Shift} = \text{ActualRT} - \text{ExpRT}$$

where

Shift	Time shift of the compound
ActualRT	Actual retention time of the compound
ExpRT	Expected retention time of the compound

$$\text{NewExpRT} = \text{ExpRT} + \left(\text{Shift} * \frac{\text{RTUpdate}}{100} \right)$$

where

NewExpRT	Corrected expected retention time of the compound
ExpRT	Expected retention time of the compound
Shift	Time shift of the compound
RTUpdate	Weighting factor (RT update factor (%)) of the compound

Updated Retention Times with Time Reference Compounds

You can use RT updates with or without *time references*. The shift and the corrected retention times of the time reference compounds themselves are calculated the same way as for any compound, using the *RT Update* function:

$$\text{NewExpRT}_{\text{Ref}} = \text{ExpRT}_{\text{Ref}} + \left(\text{Shift}_{\text{Ref}} * \frac{\text{RTUpdate}_{\text{Ref}}}{100} \right)$$

where

NewExpRT _{Ref}	Corrected expected retention time of the reference compound
ExpRT _{Ref}	Expected retention time of the reference compound
Shift _{Ref}	Time shift of the reference compound
RTUpdate _{Ref}	Weighting factor (RT update factor (%)) of the reference compound

The corrected retention time of a compound associated with a time reference is calculated using an additional factor:

$$\text{NewExpRT} = \text{ExpRT} + \left(\text{Shift}_{\text{Ref}} * \frac{\text{RTUpdate}}{100} * \text{Factor} \right)$$

where

NewExpRT	Corrected expected retention time of the compound
ExpRT	Expected retention time of the compound
Shift _{Ref}	Time shift of the reference compound
RTUpdate	Weighting factor (RT update factor [%]) of the compound
Factor	Factor for compounds with associated time reference compound (Factor)

In case of *timed groups*, the expected retention time, range start time, or range end time are only updated if you use time references or relative retention times.

Example: Retention Time Updates with RRT

If you automatically update the retention times, and also use relative retention times, the values are updated as follows:

- The expected RT of the RRT reference compound is calculated as shown in the equation for NewExpRT (see ["Calculations for Updated Retention Times"](#) on page 94)
- The expected RT of the associated compound is adjusted to keep the RRT values constant, as shown in the equation for Expected RT_{compound} (see ["Calculations for Relative Retention Times \(RRT\)"](#) on page 90)
- The RT start and RT stop times of a timed group are adjusted to keep the RRT values constant, using the same equation as for the expected RT_{compound}.

For example, consider a processing method with 3 compounds and a Timed Group, where RT Update and RRT are used:

Compound Table			Qualifier Setup		General				
#	Type	Name	Signal	Is RRT ref.	Associated RRT reference	RT update	RT update factor (%)	Exp. RT (min)	RRT
1		TimedG	DAD1A	☐	RRT Ref (DAD1A)	After each run		0.000	
3		RRT Ref	DAD1A	☒		After each run	50.000	3.000	1.000
4		C2	DAD1A	☐	RRT Ref (DAD1A)	After each run		6.000	2.000
2		C3	DAD1A	☐	RRT Ref (DAD1A)	After each run		1.500	0.500

Figure 80 Compound parameters

Timed Group Parameters				
☐ Include identified peaks				
☐ Quantify each peak individually				
start	RRT start	RT stop	RRT stop	
3.000	1.00	6.000	2.00	

Figure 81 Timed group parameters

After processing an injection, the 3 compounds were found at the following retention times:

- RRT ref: 4.000 min***
- C2: 8.000 min
- C3: 2.000 min

As a result, the expected RT of the RRT reference compound is corrected as follows:

$$\text{NewExpRT} = \text{ExpRT} + \left(\text{Shift} * \frac{\text{RTUpdate}}{100} \right)$$

$$\text{NewExpRT} = 3.000 \text{ min} + \left((4.000 \text{ min} - 3.000 \text{ min}) * \frac{50}{100} \right)$$

$$= 3.500 \text{ min}$$

The expected retention times of the other compounds as well as the start and stop times of the timed group are adjusted, so that RRT is constant.

$$\text{Expected RT}_{\text{compound}} = \text{RRT} * \text{expected RT}_{\text{reference}}$$

$$\text{Expected RT}_{\text{C2}} = 2.000 * 3.500 \text{ min} = 7.000 \text{ min}$$

$$\text{Expected RT}_{\text{C3}} = 0.500 * 3.500 \text{ min} = 1.750 \text{ min}$$

$$\text{RTStart}_{\text{TimedG}} = 1.000 * 3.500 \text{ min} = 3.500 \text{ min}$$

$$\text{RTStop}_{\text{TimedG}} = 2.000 * 3.500 \text{ min} = 7.000 \text{ min}$$

Calculation for Global Retention Time Shift

As part of method editing, you may shift all expected retention times and time ranges for timed groups at once. The new retention times are calculated as follows.

Absolute shift:

$$\text{NewExpRT} = \text{ExpRT} + \text{Shift}$$

where

NewExpRT	Corrected expected retention time of the compound
ExpRT	Expected retention time of the compound
Shift	Absolute value entered by the user

Relative shift:

$$\text{NewExpRT} = \text{ExpRT} + \left(\text{ExpRT} * \frac{\text{Shift}}{100} \right)$$

where

NewExpRT	Corrected expected retention time of the compound
ExpRT	Expected retention time of the compound
Shift	Relative value entered by the user

5

Calibration

What is Calibration?	99
Calibration Curve	100
What is a Calibration Curve	100
Response Type and Response Factor	101
Calibration Level	106
Calibration Point Weighting	108
Calibration Curve Models	111
Calibration Curve Calculation	113
Parameters for Curve Calculation	114
Linear Fit	115
Quadratic Fit	116
Logarithmic and Exponential Fits	119
Average RF fit	120
Evaluating the Calibration Curve	121
Verification of the Calibration Curve	121
Relative Residuals	121
Calibration Curve Statistics	123

This chapter contains details of the calculations used in the calibration process.

What is Calibration?

After the peaks have been integrated and identified, the next step in the quantitative analysis is the calibration. The amount and response is rarely in direct proportion to the actual mass of the sample to be analyzed. This makes the calibration with reference materials necessary. Quantitation uses peak area or height to determine the amount of a compound in a sample.

A quantitative analysis involves many steps which are briefly summarized as follows:

- Know the compound you are analyzing.
- Establish a method for analyzing samples containing a known amount of this compound, which is called the calibration sample or standard.
- Analyze the calibration sample to obtain the response due to that amount.
You may alternatively analyze a number of these standards with different amounts of the compounds of interest if your detector has a non-linear response. This process is referred to as *multi-level calibration*.

With the following calibration methods you can perform quantitation:

- Compound specific calibration (ESTD, ISTD)
- Indirect quantitation using calibration or response factor from another compound or group
- Fixed response factor (**Manual Factor**)

The ESTD calibration curves and calculations are based on measured responses (area or height) of given amounts. The ISTD calibration curves and calculations are based on relative responses and relative amounts (see ["Relative responses with ISTD"](#) on page 103).

Calibration Curve

What is a Calibration Curve

A calibration curve is a graphical presentation of the amount and response data for one compound obtained from one or more calibration samples.

Normally an aliquot of the calibration sample is injected, a signal is obtained, and the response is determined by calculating the area or height of the peak, similar to the following figure.

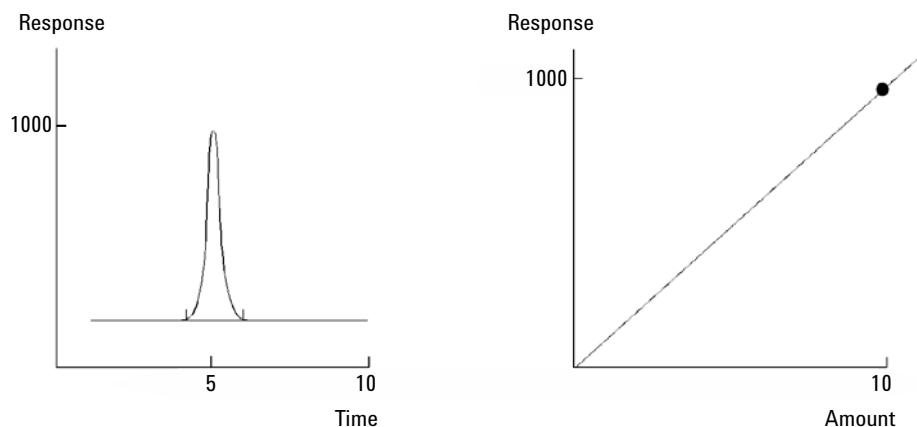


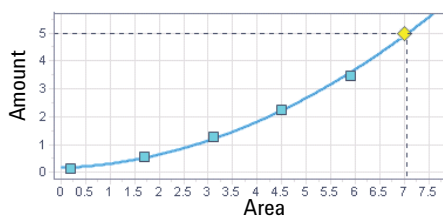
Figure 82 Calibration sample, signal, and calibration curve

Response Type and Response Factor

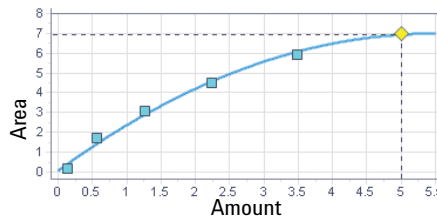
There are different settings that allow you to choose which values are plotted on the x and y axis of the calibration curve:

RF definition

The response factor (RF) is a measure of the extent to which the signal changes if a compound is detected. It is defined as the ratio of the response to the compound amount, or vice versa. In the general method settings under **RF definition**, you can switch between **Response per amount** (default) or **Amount per response**. If you change this setting, you swap the x and y axis of the calibration curve.



Amount per response



Response per amount (default)

Figure 83 Different RF definitions, Response set as Area

RF calculation:

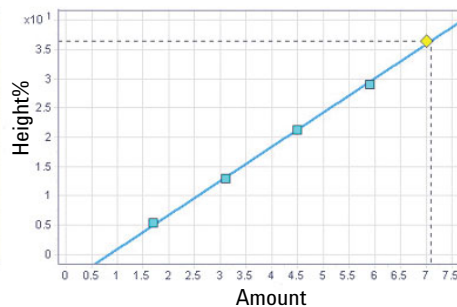
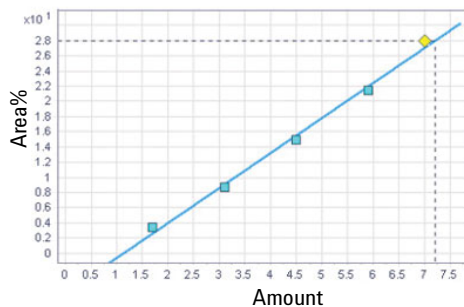
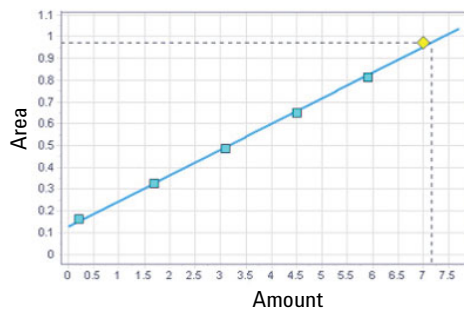
$$\text{RF} = \text{Response}/\text{Amount}$$

or

$$\text{RF} = \text{Amount}/\text{Response}$$

Type of response

The response itself can be defined as **Area**, **Area%**, **Height**, or **Height%**. You can choose the response type individually for each compound.

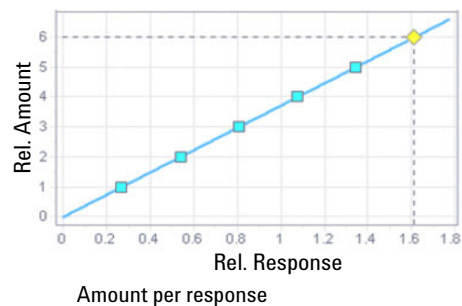
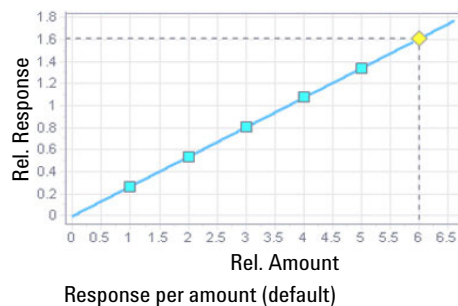


Calibration

Calibration Curve

Relative responses with ISTD

If you use internal standards (ISTDs) in your sample, relative amounts and relative responses are shown in the calibration curve. The calculation depends on the RF definition.



RF calculation:

$$RF = (\text{Response}/\text{ISTD Response}) / (\text{Amount}/\text{ISTD Amount})$$

or

$$RF = (\text{Amount}/\text{ISTD Amount}) / (\text{Response}/\text{ISTD Response})$$

Log/log curve model

If you select the curve model **log/log** for a compound, the amount and response are both plotted as logarithmic values.

You can use the **log/log** model in combination with both RF definitions, with all response types, and with internal or external standards.

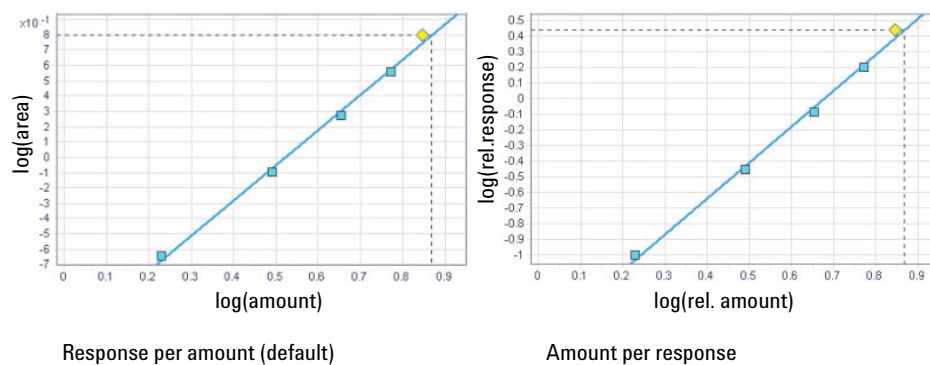


Figure 84 Example for log/log model with external or internal standards

RF calculation for the examples shown above:

$$\text{RF} = \log(\text{Response}) / \log(\text{Amount})$$

$$\text{RF} = \log(\text{Response}/\text{ISTD Response}) / \log(\text{Amount}/\text{ISTD Amount})$$

Scaled response

If you select a scaled response, the response is plotted as a calculated value while the amount is plotted without modification.

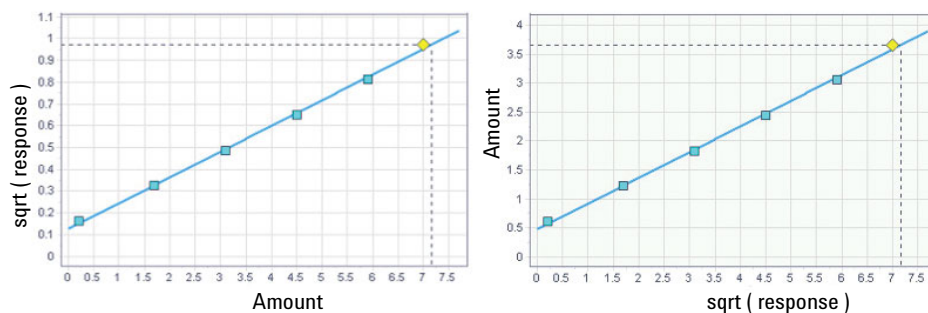


Figure 85 Examples for a scaled response using the sqrt function

RF calculation for the example shown above:

$$\text{RF} = \sqrt{\text{Response}} / \text{Amount}$$

or

$$\text{RF} = \text{Amount} / \sqrt{\text{Response}}$$

Calibration Level

There is one global number of calibration levels per processing method for all compounds. The number of calibration levels defines how many points (amount, response) are used to calculate the calibration curve. You define each level by processing the corresponding calibration standard. For each compound, the calibration curve shows the calibration points that have been used to calculate the averages.

When re-running a calibration standard, or when processing further calibration standards of a given calibration level, the calibration point for that level is updated. If you chose to use average values in the **Curve calculation** settings, the calibration point is updated by the average value of the new measured point and all already existing value(s) (see ["Modes of Using Individual Points"](#) on page 106),

The collection of calibration points can be controlled by the **Run Type** as follows:

- No selection: A new point will be added to the calibration curve.
- **Clear all Calibration:** All calibration points for all calibration levels are deleted before the new calibration points are saved.
- **Clear Calibration at Level:** All calibration points for the given calibration level are deleted before the new calibration points are saved.

Reprocessing the same calibration standard injection multiple times will update the same calibration point in the curve and not add new points.

Modes of Using Individual Points

You can choose per processing method how calibration points are used for calculating the calibration curve. The following modes are available:

- **From average per level:** Amounts and responses of all calibration points contributing to a level will be averaged and used in the algorithm to calculate the best calibration curve.
- **From individual calibration points:** All amounts and responses of the individual calibration points will be used directly to determine the calibration curve.

Average

The average from all calibration runs is calculated using the following formula:

$$\text{Response} = \frac{((n - 1) * \text{Response}) + \text{MeasResponse}}{n}$$

where

n Number of calibration points

MeasResponse	Measurement response
1	1
2	2
3	3
4	4
5	5
6	6
7	7
8	8
9	9
10	10
11	11
12	12
13	13
14	14
15	15
16	16
17	17
18	18
19	19
20	20
21	21
22	22
23	23
24	24
25	25
26	26
27	27
28	28
29	29
30	30
31	31
32	32
33	33
34	34
35	35
36	36
37	37
38	38
39	39
40	40
41	41
42	42
43	43
44	44
45	45
46	46
47	47
48	48
49	49
50	50
51	51
52	52
53	53
54	54
55	55
56	56
57	57
58	58
59	59
60	60
61	61
62	62
63	63
64	64
65	65
66	66
67	67
68	68
69	69
70	70
71	71
72	72
73	73
74	74
75	75
76	76
77	77
78	78
79	79
80	80
81	81
82	82
83	83
84	84
85	85
86	86
87	87
88	88
89	89
90	90
91	91
92	92
93	93
94	94
95	95
96	96
97	97
98	98
99	99
100	100

Bracketed Calibration

With bracketed calibrations, the samples are bracketed by pre-sample and post-sample calibrations. The calibration standards between opening and closing brackets are processed first, and a calibration curve is calculated. This curve is then used to calculate the samples in this bracket.

For all bracketing modes except **Custom**, a **Clear all calibration** operation is automatically performed for all opening brackets. For the bracketing mode **Custom**, you can set the Run Type to **Clear all calibration** as required.

Bracketing is configured in the **Injection List** window. There are different bracketing modes:

- Overall

The calibration curve is calculated with all calibration standards in the sequence, starting with the first one and finishing with the last one. All samples are reprocessed *after* the calibration curve has been calculated.

- **Non overlap**

You must have at least three sets of standards in your sequence, and at least two standards in the middle. The standards in the middle of the sequence are each used in one single bracket only.

If there are more than two standards in the middle, they will be divided and allocated to the preceding and subsequent bracket. With uneven numbers of standards in the middle, the extra standard is allocated to the preceding bracket.

- **Overlap**

You must have at least three sets of standards in your sequence. Standards in the middle of the sequence are used in two brackets (the preceding and the subsequent bracket).

- Custom

Create brackets as required. In the **Run type** column, you can choose for each calibration standard individually which calibration levels shall be cleared. If you do not choose any run type, a bracket will be averaged with its predecessor.

Calibration Point Weighting

To compensate for the variance of the response at different calibration amounts, you can specify the relative weighting (or importance) of the various calibration points used to generate the curve.

The parameter that controls the weighting is **Weighting Method**. The default weight is equal weight for all levels and the maximum weight for each curve is normalized to 1.

The following weighting factors are available:

None

All calibration points have equal weight.

$$wt = 1$$

where

wt Calibration level weighting factor

1/Amount

A calibration point is weighted by the factor 1/Amount, normalized to the smallest amount so that the largest weighting factor is 1. If the origin is included, it is assigned the mean of the weightings of the other calibration points.

$$wt = \frac{\text{Minimum(Amounts)}}{\text{CurrentAmount}}$$

where

Current Amount	Level amount
Minimum (Amounts)	Lowest amount across all points (levels) used for the calibration curve
wt	Calibration level weighting factor

1/Amount squared

A calibration point is weighted by the factor 1/Amount², normalized to the smallest amount so that the largest weighting factor is 1. Quadratic calibration point weightings can be used, for example, to adjust for a spread in calibration points. It makes sure that calibration points closer to the origin, which can normally be measured more accurately, get a higher weight than calibration points further away from the origin, which may be spread.

$$wt = \frac{\text{Minimum(Amounts)}^2}{\text{CurrentAmount}^2}$$

where

Current Amount	Level amount
Minimum (Amounts)	Lowest amount across all points (levels) used for the calibration curve
wt	Calibration level weighting factor

1/Response

A calibration point is weighted by the factor 1/Response, normalized to the smallest response so that the largest weighting factor is 1. If the origin is included, it is assigned the mean of the weightings of the other calibration points.

$$wt = \frac{\text{Minimum(Responses)}}{\text{CurrentResponse}}$$

where

Current Response	Level response
Minimum (Responses)	Lowest response across all points (levels) used for the calibration curve
wt	Calibration level weighting factor

1/Response squared

A calibration point is weighted by the factor 1/Response², normalized to the smallest response so that the largest weighting factor is 1. Quadratic calibration point weightings can be used, for example, to adjust for a spread in calibration points. It makes sure that calibration points closer to the origin, which can normally be measured more accurately, get a higher weight than calibration points further away from the origin, which may be spread.

$$wt = \frac{\text{Minimum(Responses)}^2}{\text{CurrentResponse}^2}$$

where

Current Response	Level response
Minimum (Amounts)	Lowest response across all points (levels) used for the calibration curve
wt	Calibration level weighting factor

Calibration Curve Models

OpenLab CDS can calculate the calibration according to different models. The following models are supported (see [“Calibration Curve Calculation”](#) on page 113):

- Linear fit (see [“Linear Fit”](#) on page 115)
- Quadratic fit (see [“Quadratic Fit”](#) on page 116)
- Logarithmic and exponential fit (see [“Logarithmic and exponential fits”](#) on page 119)
- Average RF fit (see [“Average RF fit”](#) on page 120)

You can set the calibration curve model individually for each calibrated compound.

Origin Handling

The application can consider the origin of the graph in different ways when calculating the calibration curve. You can set this parameter independently for each compound. Depending on the curve type, only specific origin handling options are available (for example, you cannot force the curve through the origin with a logarithmic calibration curve).

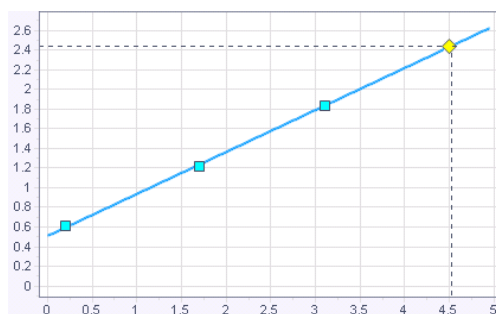


Figure 86 Ignore origin

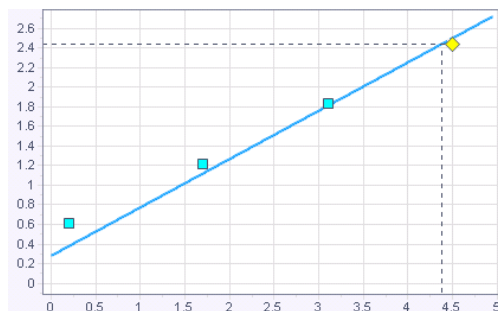


Figure 87 Include origin into the calculation

With the **Include** option, a point with amount=0 and response=0 is added to the calibration levels.

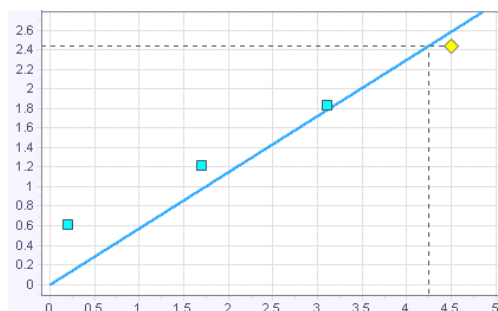


Figure 88 Force calibration curve through the origin

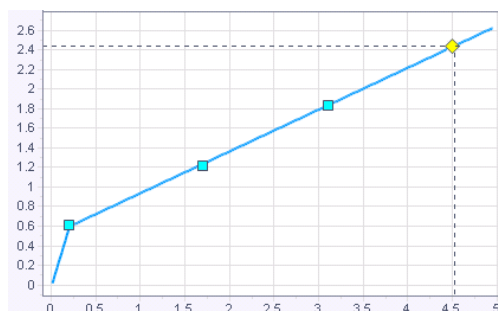


Figure 89 Connect the origin (not included in the calculation)

Calibration Curve Calculation

The optimal calibration curve is calculated by matching the curve to the calibration points. The curve calculation is based on a least squares fit (LSQ), which minimizes the sum of the residual squares. The curve type is applied to weighted calibration points. The calculation depends on the definition of the response factor (RF definition, see “Response Type and Response Factor” on page 101).

With RF defined as **Response per amount**:

$$\Sigma(wt * (CalPointArea - CalculatedArea)^2) = \min$$

where

Σ	Sum over the calibration points (levels)
CalculatedArea	The area read from the curve at calibration level amount
CalPointArea	Calibration level area
wt	Calibration point weighting factor

$$\Sigma(wt * (CalPointHeight - CalculatedHeight)^2) = \min$$

where

Σ	Sum over the calibration points (levels)
CalculatedHeight	The height read from the curve at calibration level amount
CalPointHeight	Calibration level height
wt	Calibration point weighting factor

With RF defined as **Amount per response**:

$$\sum(wt * (CalPointAmount - CalculatedAmount)^2) = \min$$

where

Σ	Sum over the calibration points (levels)
CalculatedAmount	The amount read from the curve at calibration level area or height
CalPointAmount	Calibration level amount
wt	Calibration point weighting factor

Parameters for Curve Calculation

The curve calculations all use the following parameters:

a, b, c	Curve coefficients
x	With Response per amount : Amount (ESTD), or amount ratio (ISTD) With Amount per response : Area, area%, height, or height% (ESTD) Area ratio or height ratio (ISTD)
y	With Response per amount : Area, area%, height, or height% (ESTD) Area ratio or height ratio (ISTD) With Amount per response : Amount (ESTD), or amount ratio (ISTD)
wt	Calibration level weighting factor

Linear Fit

The curve calculation is based on the least squares fit (see "Calibration Curve Calculation" on page 113).

Curve formula:

$$y = a * x + b$$

where

a Slope

b Y-intercept

Calculation of curve coefficients:

$$b = \frac{\sum(x^2 * wt) * \sum(y * wt) - \sum(x * y * wt) * \sum(x * wt)}{\sum(wt) * \sum(x^2 * wt) - \sum(x * wt)^2}$$

$$a = \frac{\sum(wt) * \sum(x * y * wt) - \sum(x * wt) * \sum(y * wt)}{\sum(wt) * \sum(x^2 * wt) - \sum(x * wt)^2}$$

At least two calibration points are required for a linear fit.

Include origin

If the origin is included, the point (0,0) is added to the other points and weighted by the mean value of the weights of the other points, that is, the $\sum(wt)$ term is increased by the mean value of the weights of the other points.

Force origin

If the force origin option is selected the curve formula is as follows:

$$y = a * x$$

where

a Slope

Calculation of curve coefficient:

$$a = \frac{\sum(x * y * wt)}{\sum(x^2 * wt)}$$

Only one calibration level is required when the origin is included or forced.

Quadratic Fit

Quadratic curve formula:

$$y = (a * x^2) + (b * x) + c$$

At least three calibration points are required for the quadratic fit. Two points are required if the origin is included or forced.

Calculation of coefficients for quadratic fit

The coefficients result from the below simultaneous linear equations. Crout's algorithm is used to solve the corresponding normal matrix equation ($A^T A x = A^T y$). In the given formula, sums are abbreviated as:

$$\begin{aligned} W &= \sum(wt) \\ XW &= \sum(x * wt) \\ X2W &= \sum(x^2 * wt) \\ X3W &= \sum(x^3 * wt) \\ X4W &= \sum(x^4 * wt) \\ YW &= \sum(y * wt) \\ XYW &= \sum(x * y * wt) \\ X2YW &= \sum(x^2 * y * wt) \end{aligned}$$

In order to avoid overflow, the x-values are normalized before entering calculation:

$$\begin{aligned} \text{Norm} &= \sum(x) \\ x &= x / \text{Norm} \end{aligned}$$

Normal equations for quadratic curve:

$$\begin{aligned} \sum(wt) * c + \sum(x * wt) * b + \sum(x^2 * wt) * a &= \sum(y * wt) \\ \sum(x * wt) * c + \sum(x^2 * wt) * b + \sum(x^3 * wt) * a &= \sum(x * y * wt) \\ \sum(x^2 * wt) * c + \sum(x^3 * wt) * b + \sum(x^4 * wt) * a &= \sum(x^2 * y * wt) \end{aligned}$$

Or written as matrix equation:

$$\begin{bmatrix} W & XW & X2W \\ XW & X2W & X3W \\ X2W & X3W & X4W \end{bmatrix} * \begin{bmatrix} c \\ b \\ a \end{bmatrix} = \begin{bmatrix} YW \\ XYW \\ X2YW \end{bmatrix}$$

Crout's decomposition:

$$\begin{vmatrix} W & XW & X2W \\ XW & X2W & X3W \\ X2W & X3W & X4W \end{vmatrix} = \begin{vmatrix} L11 & & \\ L21 & L22 & \\ L31 & L32 & L33 \end{vmatrix} * \begin{vmatrix} 1 & U12 & U13 \\ & 1 & U23 \\ & & 1 \end{vmatrix}$$

With value abbreviations:

$$L11 = W$$

$$U12 = \frac{XW}{L11}$$

$$L21 = XW$$

$$U13 = \frac{X2W}{L11}$$

$$L31 = X2W$$

$$L22 = X2W - L21 * U12$$

$$U23 = \frac{X3W - L21 * U13}{L22}$$

$$L32 = X3W - L31 * U12$$

$$L33 = X4W - (L31 * U13) - (L32 * U23)$$

$$z0 = \frac{YW}{L11}$$

$$z1 = \frac{XYW - (L21 * z0)}{L22}$$

$$z2 = \frac{X2YW - (L31 * z0) - (L32 * z1)}{L33}$$

$$a' = z2$$

$$b' = z1 - (U23 * a')$$

$$c' = z0 - (U12 * b') - (U13 * a')$$

Finally, the normalization must be reversed:

$$c = c'$$

$$b = \frac{b'}{\text{Norm}}$$

$$a = \frac{a'}{\text{Norm}^2}$$

Force Origin

If the force origin option is selected, the offset term a is set to zero when creating the normal equations.

$$\begin{vmatrix} X2W & X3W \\ X3W & X4W \end{vmatrix} \begin{vmatrix} b \\ a \end{vmatrix} = \begin{vmatrix} XYW \\ X2YW \end{vmatrix}$$

$$L11 = X2W$$

$$U12 = \frac{X3W}{L11}$$

$$L21 = X3W$$

$$L22 = X4W - (L21 * U12)$$

$$z0 = \frac{XYW}{L11}$$

$$z1 = \frac{X2YW - (L21 * z0)}{L22}$$

$$a' = z1$$

$$b' = z0 - (U12 * a')$$

$$b = \frac{b'}{\text{Norm}}$$

$$a = \frac{a'}{\text{Norm}^2}$$

Include origin

If the origin is included, the point (0,0) is added to the other points and weighted by the mean value of the weights of the other points, that is, the $\Sigma(\text{wt})$ term is increased by the mean value of the weights of the other points.

Logarithmic and Exponential Fits

Logarithmic and exponential fits

To calculate the exponential and logarithmic fit, the amount or response scales are transformed using the $\ln$ function. The linear curve fit and the weight factors are applied to the transformed data, and the curve is calculated on the transformed data.

The **Include origin** and **Force origin** options are not valid due to the singularity of the $\ln$ function at the origin.

Logarithmic

Curve formula:

$$y = a * \ln(x) + b$$

Transformations: The x scale is transformed.

$$x' = \ln(x); y' = y$$

$$y' = a * x' + b$$

Exponential

Curve formula:

$$y = b * e^{a * x}$$

Transformations: The y scale is transformed.

$$x' = x; y' = \ln(y)$$

$$y' = \ln(b) + a * x'$$

Log/log fit

To calculate the log/log fit, both amount and response scales are transformed using the log function. The linear curve fit and the weight factors are applied to the transformed data, and the curve is calculated on the transformed data.

The **Include origin** and **Force origin** options are not valid due to the singularity of the log function at the origin.

Curve formula:

$$\log(y) = a * \log(x) + b$$

Transformations: The x and y scales are transformed.

$$x' = \log(x); y' = \log(y)$$

$$y' = a * x' + b$$

Average RF fit

Average RF formula:

$$y = a * x$$

where

$$a \qquad \qquad \text{Slope}$$

The application first calculates the response factor (RF) for each calibration sample. Then, all response factors are averaged. The Average RF is indicated in the calibration curve legend and corresponds to the slope of the line ("[Response Type and Response Factor](#)" on page 101).

The y-intercept is always zero, as the force origin option is automatically used. The weighting method is automatically set to **None**.

With the Average RF fit, r and R² are not relevant and not calculated.

Evaluating the Calibration Curve

The quality of the fit of the calibration curve to the calibration levels, and the presence of outliers (measurements are at a long distance from the curve) can be evaluated using statistical calculations. The calibration curve calculation provides a correlation coefficient and a relative standard deviation for each curve, as well as a relative residual value for each calibration level.

Verification of the Calibration Curve

After calculations the calibration curves are verified and warnings are set if:

- there are not enough calibration points for the curve calculation
- the curve slope gets zero or negative
- the slope is infinite
- the calibration curve cannot be calculated (for example numeric overflows)

Relative Residuals

Residual is a measure of the calibration point distance from the calculated curve:

$$\text{Residual} = y_i - Y_i$$

where

y_i	Measured response (area or height) or amount, depending on the calibration mode.
Y_i	Predicted response or amount for level i (calculated using the curve)

The relative residual is calculated for each calibration level using the following formula:

$$\text{Rel Residual} = \frac{\text{Residual}}{Y_i} = \frac{(y_i - Y_i)}{Y_i}$$

where

y_i Measured response (area or height) or amount

Y_i Predicted response or amount for level i (calculated using the curve)

The relative residual is frequently reported in % units (**RelResidual%**). In that case the **RelResidual** needs to be multiplied by 100.

Calibration Curve Statistics

Calibration curve statistics

The calibration curve calculation provides for each curve the correlation coefficient, coefficient of determination and residual standard deviation figures.

Correlation Coefficient

The correlation coefficient (r) gives a measure of the fit of the calibration curve between the data points. It is calculated using the following equation:

$$r = \frac{\sum (y_i - \bar{y}) * (Y_i - \bar{Y}) * wt_i}{(\sum (y_i - \bar{y})^2 * wt_i) * \sum (Y_i - \bar{Y})^2 * wt_i)^{\frac{1}{2}}}$$

where

r	Correlation coefficient
wt _i	Weight of the data point
$\bar{y}$	Mean values of the measured responses or amounts If the calibration curve is forced through the origin (Origin=Force in the processing method), OpenLab CDS calculates the uncentered determination coefficient. In this case, $\bar{y}$ is omitted.
y _i	Measured response (Area, AreaRatio (ISTD method), Height or HeightRatio (ISTD method)) or amount (Amount, AmountRatio (ISTD Method)), depending on calibration mode
$\bar{Y}$	Mean values of the predicted responses or amounts
Y _i	Predicted response or amount (using the calibration curve)

$\bar{y}$ and $\bar{Y}$ are mean values of the measured and predicted responses or amounts, calculated as follows:

$$\bar{y} = \frac{\sum(y_i * wt_i)}{\sum(wt_i)}$$

where

wt_i	Weight of the data point
$\bar{y}$	Mean values of the measured responses or amounts
y_i	Measured response (Area, AreaRatio (ISTD method), Height or HeightRatio (ISTD method)) or amount (Amount, AmountRatio (ISTD Method)), depending on calibration mode

and

$$\bar{Y} = \frac{\sum(Y_i * wt_i)}{\sum(wt_i)}$$

where

wt_i	Weight of the data point
$\bar{Y}$	Mean values of the predicted responses or amounts
Y_i	Predicted response or amount (using the calibration curve)

For **Forced Origin** it is assumed that the points are centered on zero (mirrored to third quadrant) and the mean values are substituted with zero. 3rd party calculation programs may use a different approach, which will lead to slightly different results.

The correlation coefficient is 1 for a perfect fit. It reduces as the individual or averaged calibration points deviate from the regression curve. Typical values are between 0.99 and 1. The correlation coefficient is not a direct measure for precision of the analytical method, but low values indicate low precision.

Determination coefficient

The determination coefficient (R^2) is calculated as follows:

$$R^2 = 1 - \frac{\sum (y_i - Y_i)^2}{\sum (y_i - \bar{y})^2}$$

where

R^2	Determination coefficient
$\bar{y}$	Mean values of the measured responses or amounts If the calibration curve is forced through the origin (Origin=Force in the processing method), OpenLab CDS calculates the uncentered determination coefficient. In this case, $\bar{y}$ is omitted.
y_i	Measured response or amount. Response can be area (Area, Area%, or AreaRatio (ISTD method)) or height (Height, Height%, or HeightRatio (ISTD method)). Amount can be absolute amount or AmountRatio (ISTD method). The type of value depends on the calibration mode.
Y_i	Predicted response or amount (using the calibration curve)

Residual standard deviation

The residual standard deviation (sometimes referred to as the mean square error) is calculated using the following formula:

$$\text{ResidualStdDev} = \sqrt{\frac{\sum (y_i - Y_i)^2}{(n - d)}}$$

where

$d = 3$	Degree of freedom for a quadratic curve, no forced origin
$d = 2$	Degree of freedom for a quadratic curve with forced origin, or Degree of freedom for a linear curve, no forced origin
$d = 1$	Degree of freedom for a linear curve with forced origin
ResidualStdDev	Residual standard deviation
y_i	Measured response (Area, AreaRatio (ISTD method), Height or HeightRatio (ISTD method)) or amount (Amount, AmountRatio (ISTD Method)), depending on calibration mode
Y_i	Predicted response or amount (using the calibration curve)
n	number of calibration points

For **Include origin** calibration curve types, the origin (0,0) is included as a regular point in the calculation and counted by n.

The y values are not weighted.

The residual standard deviation gives a more sensitive measure of the curve quality than does the correlation coefficient. For a perfect fit, the residual standard deviation is zero. With increasing residual standard deviation values, the calibration points get further away from the curve.

Standard deviation

The standard deviation is calculated with the formula for the *population standard deviation*:

$$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2}$$

where

σ	Standard deviation
N	Number of samples
x_i	Measured value response or amount. For the curve model <i>Average RF</i> , it is the response factor RF of a compound in a single sample.
μ	Mean value. For the curve model <i>Average RF</i> , it is the average response factor of a compound in all samples.

NOTE

For the curve model **Average RF**: Due to the normally rather small population (number of calibration points), this formula is used instead of the sample population standard deviation (with N-1 as denominator).

Relative standard deviation

The relative standard deviation is calculated as follows:

$$RSD = 100 \cdot \frac{\sigma}{\mu}$$

where

RSD	Relative standard deviation
σ	Standard deviation
μ	Mean value

6

Quantitation

What is Quantitation?	128
Quantitation Calculations	128
Correction Factors	129
Multipliers	129
Dilution Factor	129
Concentration and Mass%	130
Area% and Height%	131
Quantitation of Calibrated Compounds	132
ESTD Calculation	132
ISTD Calculation	134
Quantitation of Uncalibrated Compounds	137
Indirect Quantitation Using a Calibrated Compound	137
Quantitation Using a Manual Factor	138
Quantitation of Not Identified Peaks	140
Quantify Not Identified Peaks Using a Fixed Response Factor	140
Quantify Not Identified Peaks Using a Calibrated Compound	140
Normalization	141
Calculate normalized concentrations	142
Calculate normalized amounts	143
Quantitation of groups	144
Definition of a timed group	144
Quantify a timed group	146
Definition of a named group	150
Quantify a named group	151

This chapter describes how compounds are quantified, and explains the calculations used in quantitation.

What is Quantitation?

After the peaks have been integrated and identified, the next step in the analysis is quantitation. Quantitation uses peak area or height to determine the amount of a compound in a sample.

A quantitative analysis involves many steps which are briefly summarized as follows:

- Analyze the sample containing an unknown amount of the compound to obtain the response due to the unknown amount.
- Compare the response of the unknown amount to the response of the known amount to determine how much of the compound is present.

To obtain a valid comparison for the unknown sample response to that of the known sample, the data must be acquired and processed under identical conditions.

Quantitation Calculations

OpenLab CDS offers the following calculation procedures for determining the amount of each component present in a mixture:

- Area or Height Percent (Area% or Height%)
- Quantitation using a Manual Factor
- External standard (ESTD)
- Internal standard (ISTD)
- Indirect Quantitation using a calibrated compound

The calculations used to determine the concentration of a compound in an unknown sample depend on the type of quantitation. Each calculation procedure uses the peak area or height for the calculation and produces a different type of analysis.

Correction Factors

The quantitation calculations use different correction factors, the *multiplier* (compound or injection multiplier), and the *dilution factor*. These factors are used in the calibration procedures to compensate for variations in detector response to different sample components, concentrations, sample dilutions, sample amounts, compound purities, and for converting units.

Multipliers

The multipliers are used in each calculation formula to multiply the result for each compound. A multiplier may be used to convert units to express concentrations, or to correct the concentration and thus compensate for different purities of the standard compounds.

Multipliers are set at the injection level (injection list or sequence table) and at the compound level (calibration table, part of the processing method). In OpenLab CDS, you can configure up to 5 injection multipliers and 1 compound multiplier.

The multiplier for a known compound is:

$$\text{Multiplier} = \text{Compound Multiplier} * \text{Injection Multiplier 1} * \text{Injection Multiplier 2} * \dots$$

Dilution Factor

The dilution factor is a number by which the amount is multiplied or divided to calculate the concentration (see concentration). The dilution factors are set at injection level (**Dil. factor** columns in the injection list). You can use the dilution factor to change the scale of the results or correct for changes in sample composition during pre-analysis work. You can also use the dilution factor for any other purposes that require the use of a constant factor.

The sample dilution is a combination of up to 5 dilution factors:

$$\text{Sample Dilution} = \text{Dilution Factor 1} * \text{Dilution Factor 2} * \dots$$

Concentration and Mass%

The concentration is shown in the **Concentration** column in the injection results.

The calculation depends on the following settings (all settings are available in the **General** tab of the **Compounds > Calibration** node):

- **Concentration calculation:** with this setting you choose to use dilution factors either as a divisor or as another multiplier.
- **Calculate mass %:** calculate the concentration as a mass percentage (compound amount relative to sample amount).

With **Calculate mass %** cleared (this is the default setting):

$$\text{Concentration} = \text{Amount} * \text{Multipliers} * \text{Dilution Factors}$$

or

$$\text{Concentration} = \text{Amount} * \frac{\text{Multipliers}}{\text{Dilution Factors}}$$

With **Calculate mass %** selected:

The concentration is calculated as a mass percentage (compound amount relative to sample amount).

$$\text{Concentration} = \left(\frac{\text{Amount}}{\text{Sample Amount}} * 100 \right) * \text{Multipliers} * \text{Dilution Factors}$$

or

$$\text{Concentration} = \left(\frac{\text{Amount}}{\text{Sample Amount}} * 100 \right) * \frac{\text{Multipliers}}{\text{Dilution Factors}}$$

For more information on the calculation of multipliers and dilution factors, see ["Correction Factors"](#) on page 129.

Area% and Height%

The **Area%** calculation procedure reports the area of each peak in the signal as a percentage of the total area of all peaks in the signal. **Area%** does not require prior calibration and does not depend upon the amount of sample injected within the limits of the detector. No response factors are used. If all components respond equally in the detector, then **Area%** provides a suitable approximation of the relative amounts of components.

Area% is used routinely where qualitative results are of interest and to produce information to create the compound table required for other calibration procedures.

The **Height%** calculation procedure reports the height of each peak in the signal as a percentage of the total height of all peaks in the signal.

Correction factors are not applied in Area% or Height% calculation.

Quantitation of Calibrated Compounds

The external standard (ESTD), normalization, and internal standard (ISTD) calculation procedures require calibration and therefore use a compound table. The compound table specifies conversion of responses into the units you choose by the procedure you select.

ESTD Calculation

The ESTD procedure is the basic quantitation procedure in which both calibration and unknown samples are analyzed under the same conditions. The results from the unknown sample are then compared with those of the calibration sample to calculate the amount in the unknown.

The ESTD procedure uses absolute response factors unlike the ISTD procedure. The response factors are obtained from a calibration and then stored. In following sample runs, compound amounts are calculated by applying these response factors to the measured sample responses. Make sure that the sample injection size is reproducible from run to run, since there is no standard in the sample to correct for variations in injection size or sample preparation.

When preparing an ESTD analysis, the calculation of the amount of a particular compound in an unknown sample occurs in two steps:

- 1 An equation for the curve through the calibration points for this compound is calculated using the type of fit specified in the **Mode** and **Origin** settings in the compound table.
- 2 The amount of the compound in the unknown is calculated using the equation described above. This amount may appear in the report or it may be used in additional calculations called for by sample multiplier, compound multiplier, or dilution factor values before being reported.

Single-Level Calibration

In case of single-level calibration, the response factor is simply the ratio of the calibration point response and amount. If **Include origin** and **Force origin** are switched off, a warning is emitted.

The response factor RF is defined as a ratio of response and amount or vice versa (see “[RF definition](#)” on page 101). To calculate the RF, the application uses the compound amount of the calibration sample and the corresponding response of the calibration sample.

The formula for the single-level calibration calculation of the ESTD results depends on the type of response that you have set in the processing method:

$$\text{Amount} = \text{Peak Area} / \text{RF}$$

or:

$$\text{Amount} = \text{Peak Height} / \text{RF}$$

where

Amount

Amount of the compound

RF

Response factor

For details on the calculation of concentrations, see “[Concentration and Mass%](#)” on page 130.

Multi-Level Calibration

For multi-level calibration, the response factor is evaluated from the calibration curve.

ISTD Calculation

The ISTD procedure eliminates the disadvantages of the ESTD method by adding a known amount of a compound which serves as a normalizing factor. This compound, the *internal standard*, is added to both calibration and unknown samples.

The compound used as an internal standard should be similar to the calibrated compound, both chemically and in retention/migration time, but it must be chromatographically distinguishable.

Table 8 **ISTD procedure**

Advantages	Disadvantages
Sample-size variation is not critical.	The internal standard must be added to every sample.
Instrument drift is compensated by the internal standard.	
The effects of sample preparations are minimized if the chemical behavior of the ISTD and unknown are similar.	

If the ISTD procedure is used for calibrations with a non-linear characteristic, care must be taken that errors which result from the calculation principle do not cause systematic errors. In multi-level calibrations, the amount of the ISTD compound should be kept constant, i.e. the same for all levels.

In the internal standard analysis, the amount of the compound of interest is related to the amount of the internal standard component by the ratio of the responses of the two peaks.

OpenLab CDS allows up to 5 ISTD compounds.

For the ISTD calculation relative responses and relative amounts are used instead of the "raw" responses and amounts. They are calculated by dividing the response and amount of the peak of interest by the response and amount of the corresponding ISTD compound:

$$\text{Relative Response} = \text{Response} / \text{Response}_{\text{ISTD}}$$

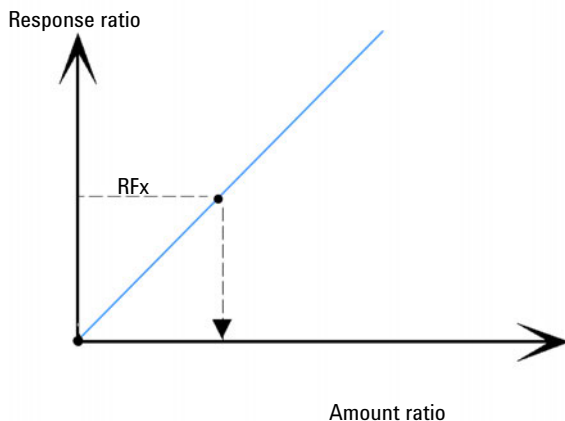
$$\text{Relative Amount} = \text{Amount} / \text{Amount}_{\text{ISTD}}$$

The response can be Area, Area%, Height, or Height% (see ["Response Type and Response Factor"](#) on page 101).

In an ISTD calibration, the calculation of the corrected amount ratio of a particular compound in an unknown sample occurs in several stages. These stages are described in the following sections.

Calibration Samples

- 1 The calibration points are constructed by calculating an amount ratio and a response ratio for each level of a particular compound in the compound table.
The amount ratio is the amount of the compound divided by the amount of the internal standard at this level.
The response ratio is the response (area or height) of the compound divided by the response of the internal standard at this level.
- 2 An equation for the curve through the calibration points is calculated using the type of curve model specified in the compound table of the processing method.



Unknown Sample

- 1 The response of the compound in the unknown sample is divided by the response of the internal standard in the unknown sample to give a response ratio for the unknown.
- 2 An amount ratio for the unknown is calculated using the curve model equation determined in "Calibration Samples" on page 135, and the actual amount of ISTD in the sample.

Single Level ISTD Calibration

In case of single-level calibration, the relative response factor (RRF) is evaluated using response and amount values from calibration samples. Depending on the RF definition in the global calibration settings, one of the following formulas applies:

With RF defined as **Response per amount**:

$$RRF = \frac{\text{Rel Response}}{\text{Rel Amount}}$$

With RF defined as **Amount per Response**:

$$RRF = \frac{\text{Rel Amount}}{\text{Rel Response}}$$

where

RRF	Relative response factor
RelResponse	Relative response
RelAmount	Relative amount

The amount and concentration are calculated according to the following formulas, using the response value from the sample measurement:

With RF defined as **Response per amount**:

$$\text{Amount} = \left(\frac{\text{Rel Response}}{RRF} \right) * \text{Amount}_{ISTD}$$

With RF defined as **Amount per Response**:

$$\text{Amount} = RRF * \text{RelResponse} * \text{Amount}_{ISTD}$$

where

Amount	Compound amount
RelResponse	Relative response
RRF	Relative response factor
Amount _{ISTD}	Amount of internal standard

For multi-level calibration, the relative response factor is evaluated from the calibration curve.

For details on calculation of concentrations, see "Concentration and Mass%" on page 130.

Quantitation of Uncalibrated Compounds

Uncalibrated compounds can be quantified either with a fixed response factor or using the calibration data of one of the calibrated compounds. Quantitation using a fixed response factor or calibrated compound data is signal-specific. In the latter case, if the calibrated compound is quantified by an ISTD method, the ISTD is used for the not identified peaks in the same way as for the calibrated compound.

Indirect Quantitation Using a Calibrated Compound

If the calibration data of a calibrated compound is to be used to quantify uncalibrated compounds, the calibrated compound is identified in the processing method (**Calibration** node, **Compound Table** tab: under **Mode**, select **Reference**). Calculations are the same as for calibrated compounds. If the reference compound is quantified by an ISTD method, the ISTD is used for the uncalibrated compound in the same way as for the reference compound.

A missing reference peak results in a zero amount of the uncalibrated compound.

Optionally, a correction factor (**Ref. correction**) can be entered to multiply the response of the peak before the amount is calculated from the response factor of the reference compound.

Quantitation Using a Manual Factor

The software allows you to quantify an identified compound that is based on a fixed response factor (**Manual Factor** column). In this case, the compound amount is calculated using the fixed response factor:

$$\text{Amount} = \text{Response} * \text{Manual Factor}$$

where

Manual Factor

Fixed response factor

Response

Response can be Area, Area%, Height, Height%, Rel. Area, or Rel. Amount (see ["Response Type and Response Factor"](#) on page 101)

For details on calculation of concentrations, see ["Concentration and Mass%"](#) on page 130.

Using a manual factor with an ISTD method

If the compound amount is quantified using the fixed response factor and ISTD, the formula is read as follows:

$$\text{Relative area} = \text{Area} / \text{Area}_{\text{ISTD}}$$

or:

$$\text{Relative height} = \text{Height} / \text{Height}_{\text{ISTD}}$$

The amount is then calculated as follows:

$$\text{Amount} = \text{Relative Area} * \text{Manual Factor} * \text{Amount}_{\text{ISTD}}$$

or:

$$\text{Amount} = \text{Relative Height} * \text{Manual Factor} * \text{Amount}_{\text{ISTD}}$$

For details on calculation of concentrations, see ["Concentration and Mass%"](#) on page 130.

Dependency of manual factor and response factor (RF)

With RF defined as **Response per amount** (default setting):

$$RF = 1 / \text{Manual Factor}$$

With RF defined as **Amount per response**:

$$RF = \text{Manual Factor}$$

For more information on the response factor, see [“Response Type and Response Factor”](#) on page 101.

Quantitation of Not Identified Peaks

Not identified peaks can be quantified using timed groups, either with a fixed response factor or with the calibration data of one of the calibrated compounds. Quantitation using a fixed response factor or calibrated compound data is signal-specific. In the latter case, if the calibrated compound is quantified by an ISTD method, the ISTD is used for the not identified peaks in the same way as for the calibrated compound.

For more information on timed groups, see ["Definition of a timed group"](#) on page 144.

Quantify Not Identified Peaks Using a Fixed Response Factor

In this case, you create a timed group with quantitation mode **Manual Factor**. The specified times ranges of the timed group include the relevant not identified peaks.

In addition, quantified peaks must be excluded. By setting the option **Quantify each peak individually**, amount and concentration of all not identified peaks are calculated using the fixed response factor.

For details on the calculation, see ["Quantify a timed group with a manual factor"](#) on page 146.

Quantify Not Identified Peaks Using a Calibrated Compound

In this case, you create a timed group with quantitation mode **Reference**. The specified times ranges of the timed group include the relevant not identified peaks. Optionally, a correction factor (**Ref. correction**) can be entered to multiply the response of the peak before the amount is calculated from the response factor of the reference compound.

In addition, quantified peaks must be excluded. By specifying the option **Quantify each peak individually**, amount and concentration of all not identified peaks are calculated using the curve reference.

Normalization

You can choose to normalize amounts in the general calibration settings of a processing method.

The normalization analysis has the same disadvantage as the Area% and Height% calculations. Any changes that affect the total peak area will affect the concentration calculation of each individual peak. The normalization analysis should only be used if all components of interest are eluted and integrated. Excluding selected peaks from a normalization analysis will change the reported results in the sample.

For detailed information on timed groups, see [“Definition of a timed group”](#) on page 144.

If individual peaks are calculated in timed groups, these individual peaks are not included a second time in the total amount; they are already included in the timed group amount.

The default normalization factor is 100.00 to create Norm% results. However, you can set a different number and unit in the method. The total amount is the sum of all calculated compound and timed group amounts, independent of the signal of the compound main peak.

You can select whether ISTD compounds are included in the calculation. If excluded (default setting), the ISTD amount is not added to the total amount, and no compound normalization amounts are calculated for the ISTDs.

For *named groups* the group amount is not included in the total amount. However, if an ISTD has been explicitly added to this named group, it will add to the group amount, even if ISTD is excluded from normalization. This can result in normalization amounts greater than 100% for the named group.

Calculate normalized concentrations

The normalized concentration of a compound refers to the concentration of the compound in relation to the sum of concentrations of all detected compounds.

The calculation depends on the status of the **Apply correction factors to normalized concentrations** check box. This check box is part of the calibration settings in the processing method.

If correction factors are applied, the **Normalized concentration** of a compound is calculated as follows, depending on the concentration calculation:

$$C_N = \frac{M_{\text{comp}} \cdot A}{\sum (M_{\text{comp}} \cdot A)} \cdot M_{\text{sample}} \cdot D_{\text{sample}} \cdot f$$

Or

$$C_N = \frac{M_{\text{comp}} \cdot A}{\sum (M_{\text{comp}} \cdot A)} \cdot \frac{M_{\text{sample}}}{D_{\text{sample}}} \cdot f$$

If correction factors are not applied, the normalized concentration is calculated as:

$$C_N = \frac{M_{\text{comp}} \cdot A}{\sum (M_{\text{comp}} \cdot A)} \cdot f$$

where

C_N	Normalized concentration
M_{comp}	Compound specific multiplier, set in the processing method
A	Amount
M_{sample}	Sample specific multipliers, set in the injection list or sequence table
D_{sample}	Sample-specific dilution factors, set in the injection list or sequence table
f	Normalization factor, defined in the processing method

Calculate normalized amounts

The equation used to calculate the **Norm amount** of a compound is:

$$\text{Compound Norm Amount} = \text{Compound Amount} * \frac{\text{Normalization}}{\text{Total Amount}}$$

where

Compound Amount	Amount of compound
Compound Norm Amount	Amount of normalization compound
Normalization	Normalization factor, defined in the processing method
Total Amount	Sum of all compound amounts and timed group amounts Named group amounts are not included in the total amount. Amounts of identified compounds in timed groups are counted twice if you have enabled Include identified peaks for the timed group.

$$\text{Group Norm Amount} = \text{Group Amount} * \frac{\text{Normalization}}{\text{Total Amount}}$$

where

Group Amount	Amount of group
Group Norm Amount	Amount of normalization group
Normalization	Normalization factor, defined in the processing method

$$\text{Peak Norm Amount} = \text{Peak Amount} * \frac{\text{Normalization}}{\text{Total Amount}}$$

where

Normalization	Normalization factor, defined in the processing method
Peak Amount	Amount of peak
Peak Norm Amount	Amount of normalization peak

Quantitation of groups

Definition of a timed group

A timed group contains one or more time regions and is defined on a specific signal. The expected retention time of the group is for sorting purposes only and can be entered manually.

The timed group corresponds to the *uncalibrated range* or *calibrated range* in OpenLab EZChrom.

The following example shows three timed groups, where Group 2 and Group 3 overlap. C1 and C2 are identified compounds. The unidentified peak at 5.689 min is evaluated in both groups. Identified peaks are only evaluated if the group parameters are set accordingly.

Groups	Time ranges	Include identified peaks?
Group 1	0.8 min - 1.4 min 2.8 min - 3.4 min	No
Group 2	3.8 min - 5.9 min	Yes
Group 3	5.4 min - 7.2 min	No

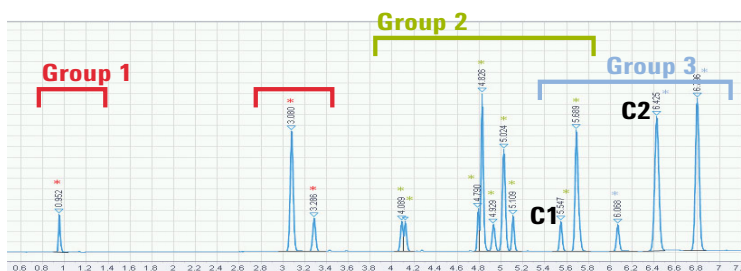


Figure 90 Example: Timed groups

NOTE

If a timed group has no time region defined, its area, height, and amount are not calculated. If a timed group has a region or regions defined, but no peak is found in this region or these regions, its area and height are equal to zero.

If you correct the retention times by using time reference compounds, the start time and stop time of the timed group are also corrected by the corresponding shifts (see [“Calculations for Time Reference Compounds”](#) on page 92).

For quantitation, the area, height, or scaled response of the group is calculated by summing the corresponding values of the individual peaks contained in the time ranges (including or excluding identified peaks, depending on the group parameters).

A timed group is calibrated and quantified using the calibration parameters for the group. Timed groups support all calibration and quantitation modes of regular compounds (**Curve, Manual Factor, Reference**). In the group parameters, you can choose to quantify all peaks individually. In this case, each peak of the group is quantified individually with the group response factor (RF).

Caution: As Timed Groups use the calibration parameters of the whole group, caution should be taken when using with nonlinear and/or calibration curves with a large offset.

Usage of a Timed Group with nonlinear calibration curves is not recommended as it will calculate potentially imprecise results due to the noncommutative nature of the calculations.

If a group calibration curve has a significant off-set (that is, a large intercept (b) value), the offset gets equally distributed across all peaks found within the time group windows. The equal distribution may not describe each compound adequately.

Conflicts may occur if a peak belongs to several groups, or if identified peaks are quantified as part of the group. For details on the conflict resolution, see [“Evaluation of timed group parameters”](#) on page 148.

Quantify a timed group

Quantify a timed group with a manual factor

In this case the amount of the group is calculated according to a fixed response factor entered manually.

ESTD:

$$\text{Group Amount} = \text{Group Response} * \text{Manual Factor}$$

where

Manual Factor	Fixed response factor
Group Response	Sum of all (scaled) responses

or ISTD:

$$\text{Group Amount} = \frac{\text{Group Response}}{\text{Response}_{\text{ISTD}}} * \text{Manual Factor} * \text{Amount}_{\text{ISTD}}$$

where

Amount _{ISTD}	Amount of the internal standard
Manual Factor	Fixed response factor
Response _{ISTD}	(Scaled) response of the internal standard

$$\text{Group Concentration} = \text{Group Amount} * \text{Multipliers} * \text{Dilution Factors}$$

or

$$\text{Group Concentration} = \text{Group Amount} * \text{Multipliers} / \text{Dilution Factors}$$

For more information on how to calculate the multipliers and dilution factors, see [“Correction Factors”](#) on page 129. For more information on how to calculate the concentration, see [“Concentration and Mass%”](#) on page 130.

Quantify a timed group with its own calibration curve

A timed group can be quantified according to its own calibration curve. All calibration options or levels are supported, except the response scaling **Response/RT**. In the ISTD mode, you must select an ISTD to use.

If you choose a scaled response value, the scaled response is the sum of the individually scaled responses.

Quantify a timed group with the curve of a reference compound

The timed group can be quantified according to the calibration curve of another single compound. The software allows you to use the response factor of a reference compound (calibration curve). In this case a correction factor (**Ref. correction**) can be entered to multiply the response before the amount is calculated from the response factor of the reference compound. In case of ISTD same ISTD is used as for reference compound.

Quantify peaks individually in a timed group

If you choose to quantify all peaks individually, the individual peak amount is calculated as:

$$\text{Peak Amount} = \text{Group Amount} * \frac{\text{Peak Response}}{\text{Group Response}}$$

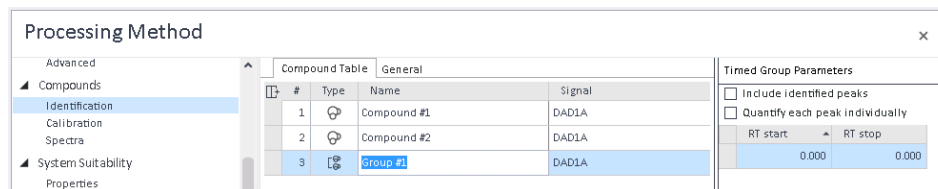
$$\text{Peak Concentration} = \text{Peak Amount} * \text{Multipliers} * \text{Dilution Factors}$$

or

$$\text{Peak Concentration} = \text{Peak Amount} * \text{Multipliers} / \text{Dilution Factors}$$

For more information on how to calculate the multipliers and dilution factors, see "[Correction Factors](#)" on page 129. For more information on how to calculate the concentration, see "[Concentration and Mass%](#)" on page 130.

Evaluation of timed group parameters



If the **Include identified peaks** check box is not selected in the timed group definition, the group results only include the results for the unknown peaks found within one or more time regions. If selected, results for identified peaks are also taken into account.

If the **Quantify each peak individually** check box is not selected in the timed group definition, the individual peaks that belong to the group are not quantified. Only area, height, and scaled response are available as individual peak results, but not the amount or concentration. If selected, each peak in the given time region is quantified with the group calibration parameters. Selecting or clearing this check box does not change the group results.

Conflict resolution

Conflicts may occur if a peak belongs to several groups, or if identified peaks are quantified as part of the group. In these cases, the following rules apply:

- If an unknown peak belongs to several groups, and calibration parameters including calibration levels are defined for both groups: The peak is reported several times, but always with the same amount. It is quantified only one time. The peak is quantified with the parameters of the group with the smallest expected retention time. If both groups have the same expected retention time, the last created group is used.
- If an identified peak is quantified as part of the group, but has its own calibration parameters defined, it is quantified with its own parameters and not with the group parameters.
- If an identified peak is quantified as part of the group and has no specific calibration parameters defined, the compound is quantified with the group parameters. The response type (area or height) of the group is used.
- If an identified peak is the *internal standard (ISTD)* of this group, and the group parameter **Include identified peaks** is set, the response of this peak will be subtracted from the timed group response.

**Quantitation
with scaled
responses**

In the **Injection Results** window, the **Scaled response** column shows the response of each quantified peak. If the group's calibration curve is not scaled, the scaled response equals to the peak area or height.

The group's scaled response is calculated as the *sum* of the scaled responses of the individual peaks that are included in the timed group. The group calibration and quantitation are based on this sum of scaled responses

Area and height of the timed group are also summed values, but they are only calculated if no scaling is applied to the response values.

Definition of a named group

The named group consists of user-selected compounds and timed groups. Each compound or timed group is identified and quantified on its own. ESTD and ISTD calculations are based on the calibration data of the individual compounds. The calculated group area, height, amount, and concentration of the group are the sum of the individual areas, heights, amounts, and concentrations. The named group itself is not calibrated. One compound can be in multiple named groups.

The expected retention time of the group is only for sorting purposes and can be entered manually.

The named group corresponds to the *Named Peaks group* in OpenLab EZChrom.

The following example shows two named groups where one compound is contained in both groups:

Groups	Group 1	Group 2
Included compounds	C1	
	C2	
		C3
	C4	C4
		C5
		C6

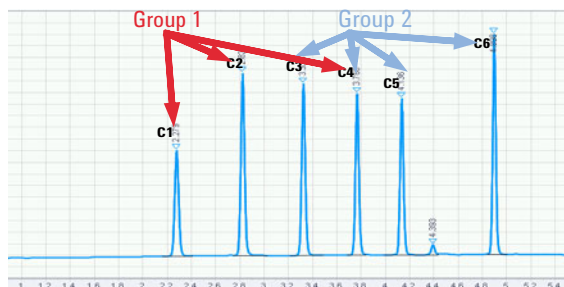


Figure 91 Example: Named groups

Quantify a named group

The results in the named group table are:

Group Area = Σ Compound Peak Areas + Σ Timed group areas

Group Height = Σ Compound Peak Heights + Σ Timed group heights

Group Amount = Σ Compound Amounts + Σ Timed group amounts

Group Concentration = Σ Compound Concentrations + Σ Timed group concentrations

If none of the compounds or timed groups in the named group have been identified, the named group will appear as "not identified" on the analysis.

What is UV spectral analysis?	153
UV impurity check	155
Noise calculation	155
Determine significant spectra	156
Background correction	156
Similarity calculation	157
Threshold calculation	158
Impurity evaluation	159
About threshold, similarity curves and sensitivity	160
Comparison with traditional purity plots	163
UV confirmation	166

This chapter describes the concepts of the impurity check and the confirmation of compound identity based on UV spectral analysis.

What is UV spectral analysis?

There are different windows and functions specific to UV spectral analysis. To view those windows and access the functions, the focused injection must contain spectral data (for example, acquired with a 3D UV system).

UV spectral analysis provides additional quality criteria for routine analytics:

- *Confirm the compound identity*

The application compares a UV spectrum with a specific UV reference spectrum. A high match factor indicates that the compounds are probably identical.

For details on the calculation, see UV confirmation ("[UV confirmation](#)" on page 166).

- *Check for UV impurities*

The application compares all UV spectra of a peak with the apex spectrum. It calculates an overall match factor, the UV Purity value. A low UV Purity value indicates that there are co-eluted peaks with a significantly different UV spectrum.

For details on the calculation, see UV impurity check ("[UV impurity check](#)" on page 155).

UV Spectral analysis processes spectral data acquired from a UV-visible diode-array detector or fluorescence detector. It adds a third dimension to your analytical data when using it with the chromatographic data (see [Figure 92](#) on page 154).

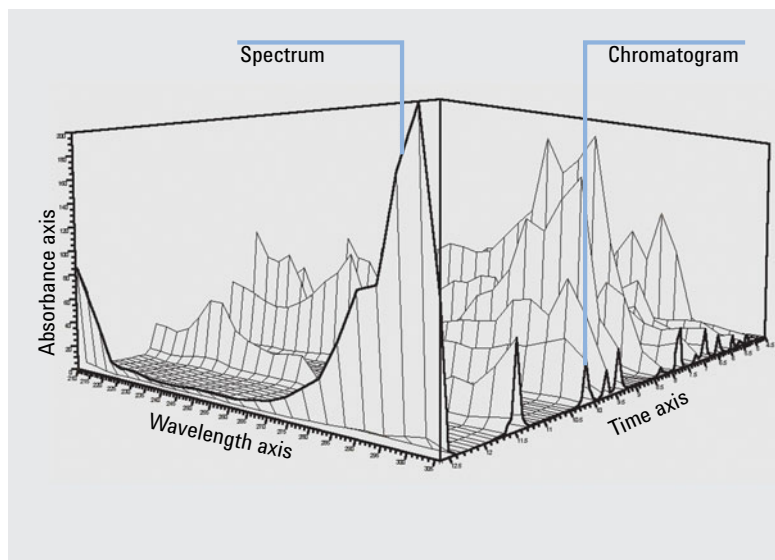


Figure 92 Spectral Information

UV impurity check

An impurity check assesses whether a peak is pure or contains impurities. This assessment is based on the comparison of spectra recorded during the elution of the peak. After applying a baseline correction, the spectrum at the peak apex is compared with all significant spectra recorded in the peak. The application calculates a match factor that characterizes the degree of similarity of the spectra.

The application performs the following steps to evaluate UV impurities:

- 1 Per peak
 - a "Noise calculation" on page 155
 - b Determine significant spectra("Determine significant spectra" on page 156)
- 2 Per spectrum:
 - a "Background correction" on page 156
 - b "Similarity calculation" on page 157
 - c "Threshold calculation" on page 158
- 3 "Impurity evaluation" on page 159

Noise calculation

As a preparation for further evaluations, the application calculates the following numbers for each peak from the spectra at baseline start and baseline end:

- Noise variance
- Noise standard deviation σ

Baseline start and end times depend on the integration. If multiple peaks are only separated by a drop line, all peaks use the same spectra for noise calculation.

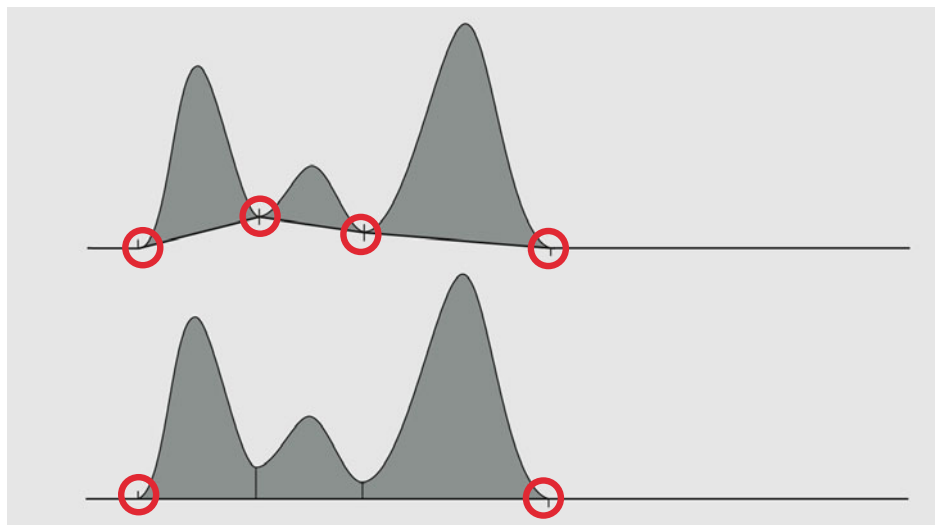


Figure 93 Baseline start and end depending on integration

Determine significant spectra

To ensure that only spectra with a significant signal are evaluated, the application filters out spectra where the response range is too small. Spectra are used for further calculations only if the following applies:

- Response range is larger than 3σ
- Response range is larger than or equal to 10 % of the apex spectrum response range. The response range for each spectrum is calculated as max - min response.

Background correction

For the baseline correction, the application evaluates the following spectra:

- Spectrum at the baseline start of the peak
- Spectrum at the baseline end of the peak

Baseline start and end times depend on the integration. If multiple peaks are only separated by a drop line, all peaks use the same spectra for background correction (see [Figure 93](#) on page 156).

A linear interpolation of the two baseline spectra is calculated. To correct each individual peak spectrum, the application subtracts the interpolation spectrum at the corresponding retention time.

Similarity calculation

The application compares each of the remaining background-corrected peak spectra with the background-corrected apex spectrum. The match factor is a value between 0 (no similarity) and 1000 (identical spectra).

$$r = \frac{\sum_{i=1}^n [(x_i - x_{av}) \cdot (y_i - y_{av})]}{\sqrt{\left[\sum_{i=1}^n (x_i - x_{av})^2 \cdot \sum_{i=1}^n (y_i - y_{av})^2 \right]}}$$

where

r	Correlation
x_i, y_i	Measured absorbances at the same wavelength from the considered data point and peak apex respectively
n	Number of wavelengths acquired per data point in time (depending of the processing method wavelength bandwidth, and acquisition spectra collection steps)
x_{av}, y_{av}	Average absorbances of the considered data point and peak apex spectra respectively

Match factor (at each data point) = $r^2 * 1000$

Threshold calculation

The reference threshold at 50% sensitivity is calculated using the following formula:

$$T = 1000 * \left(1 - 0,5 * \left(\frac{VAR_{noise}}{VAR_{peak}} + \frac{VAR_{noise}}{VAR_{target}} \right) \right)$$

where

VAR_{noise}	Calculated variance threshold of the noise spectra
VAR_{peak}	Variance of the peak spectra
VAR_{target}	Variance of the spectrum used for the comparison (peak start for data points purity calculation before the apex, and at peak end for data points purity calculation after the apex)

The threshold (T_s) used in the software depends on the sensitivity value. It is calculated with following formulas:

If the sensitivity is greater than 50%:

$$T_s = T + \frac{(1000 - T) \cdot (S - 50)}{50}$$

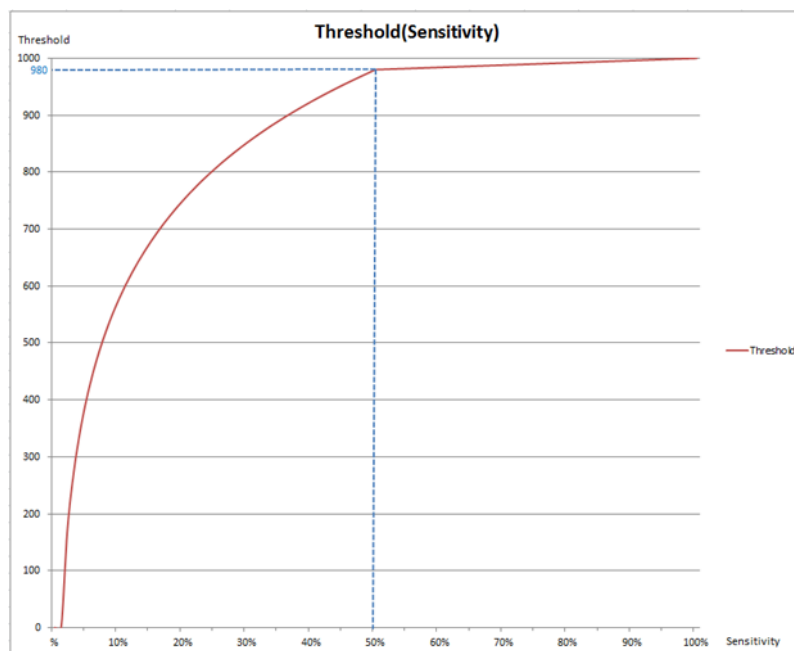
If the sensitivity is lower than or equal to 50%:

$$T_s = T \cdot \frac{\log(S)}{\log(50)}$$

where

T_s	threshold [0 – 1000] for selected sensitivity S
T	Reference threshold at 50% sensitivity
S	Sensitivity [0 - 100] %

For each individual threshold value, at each raw data point, one of these formulas is applied.



The red curve is the automatically calculated threshold curve. The automatically calculated threshold value for this curve example is 980 (which is the sensitivity value 50%).

Impurity evaluation

The single values for the similarity curve are calculated from the sensitivity-corrected threshold value and the match factor ("[Threshold calculation](#)" on page 158, "[Similarity calculation](#)" on page 157).

$$\text{Ratio} = \log((1000 - \text{Threshold}) / (1000 - \text{Match Factor}))$$

All values of a peak are shown in the *similarity curve*. You can view it in the **Peak Details** window. This similarity curve shows a distribution of positive values (pure data points), and negative values (impure data points) across the peak, where the 0 would represent the threshold limit.

About threshold, similarity curves and sensitivity

In OpenLAB CDS v2.x, the threshold is calculated automatically for every data point. The sharpness of the analysis is adjusted using the sensitivity percentage.

The sensitivity is therefore not a fixed threshold value. Modifying the sensitivity does not change the threshold value with a linear relationship, as the threshold values are not identical at every raw data point. As an example, on the figure below, if the threshold value at 1.950 min is 998, it does not necessarily mean that the threshold would be the same at 1.960 min. Every data point have its own threshold calculated, leading to a threshold curve across the peak.

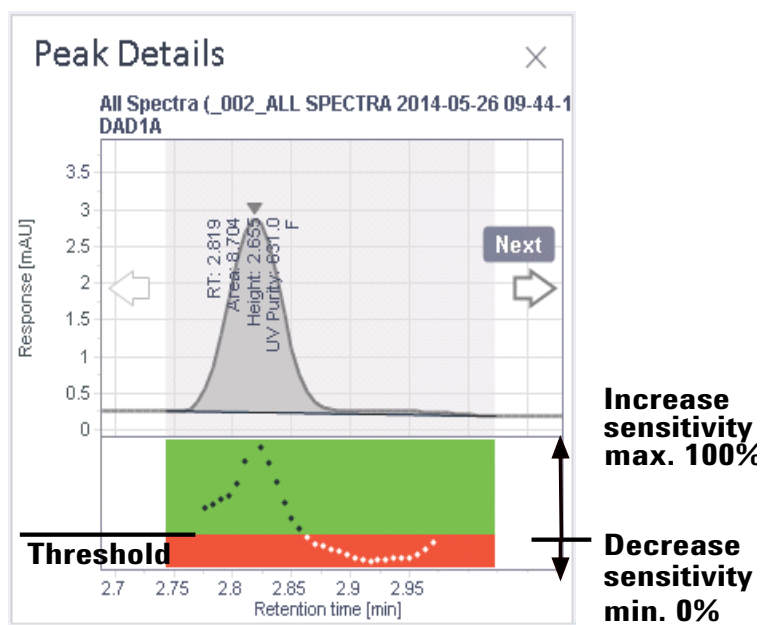


Figure 94 Similarity curve

The peak purity sensitivity is then applied to the threshold curve. In Data Analysis, a logarithmic transformation of the threshold and similarity curves is drawn below the peaks. The threshold curve is then a flat line between the pure and impure regions. The similarity curve (shown in Peak Details) is drawn using the formula as described under “Impurity evaluation” on page 159. This similarity curve is then showing a distribution of positive values (pure data points), and negative values (impure data points) across the peak, where the 0 represents the threshold limit.

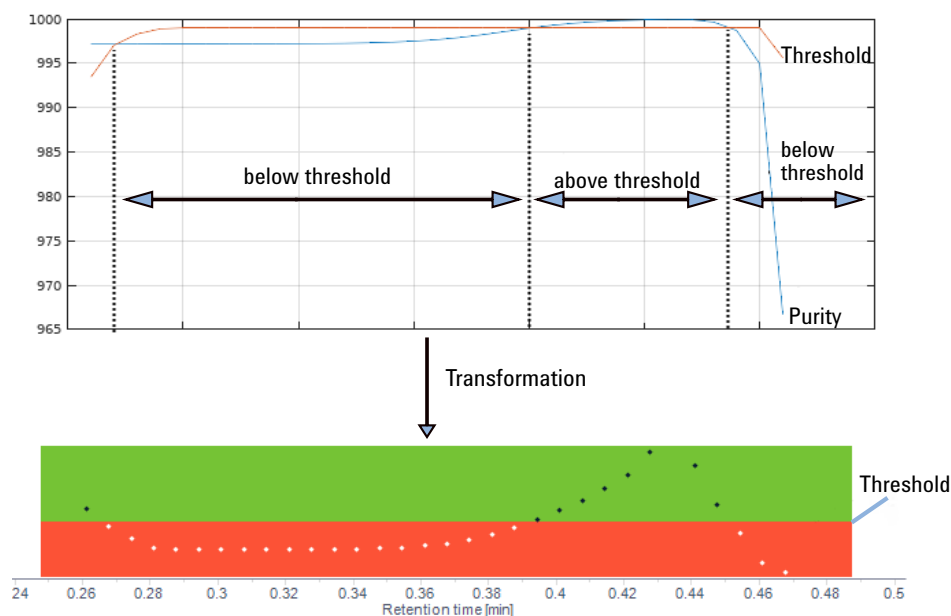


Figure 95 Before and after transformation

For example:

- Data point with a match factor of 990 and a threshold of 980:

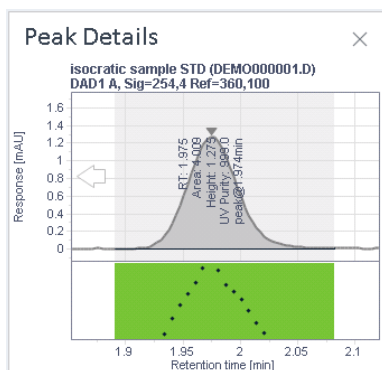
$$\text{Ratio} = \log((1000 - 980) / (1000 - 990)) = \log(2) = +0.3$$
 This raw data point meets the purity criteria.
- Data point with a match factor of 970 and a threshold of 990:

$$\text{Ratio} = \log((1000 - 990) / (1000 - 970)) = \log(0.33) = -0.48$$
 This raw data point does not meet the purity criteria.

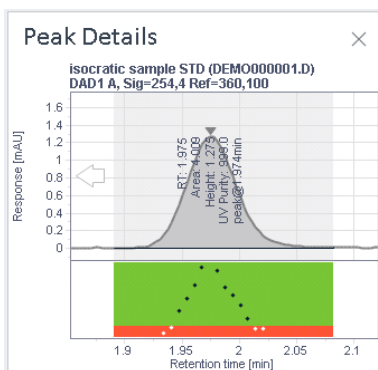
Decreasing or increasing the sensitivity means that the profile of the calculated threshold curve is changing. The full range for sensitivity is from 0 to 100 %, where the default calculated threshold is at 50 %.

Because the display of the similarity curve is nonlinear (but logarithmic), the threshold will not map to a one to one relationship from one data point to another. If the sensitivity is increased/decreased by 20 %, the threshold curve is moved up or down, and its amplitudes change as well. The raw data point thresholds are not changed by +/- 20 %.

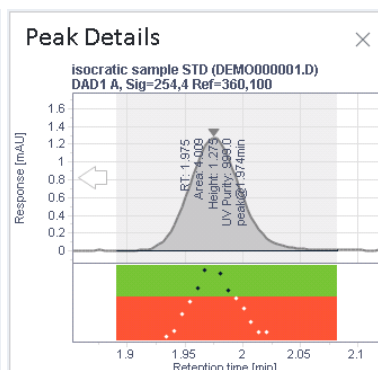
Sensitivity	Threshold
0 %	Lowest possible value = 0
$0 \leq s \leq 100$	Calculated threshold curves (with a reference curve at 50%)
100 %	Highest possible value = 1000



Low sensitivity - this peak is considered pure



Default sensitivity - this peak is considered impure



High sensitivity - this peak is considered impure

A peak is flagged as impure when one data point is under its threshold, even if the overall UV peak purity factor is close to 1000.

Comparison with traditional purity plots

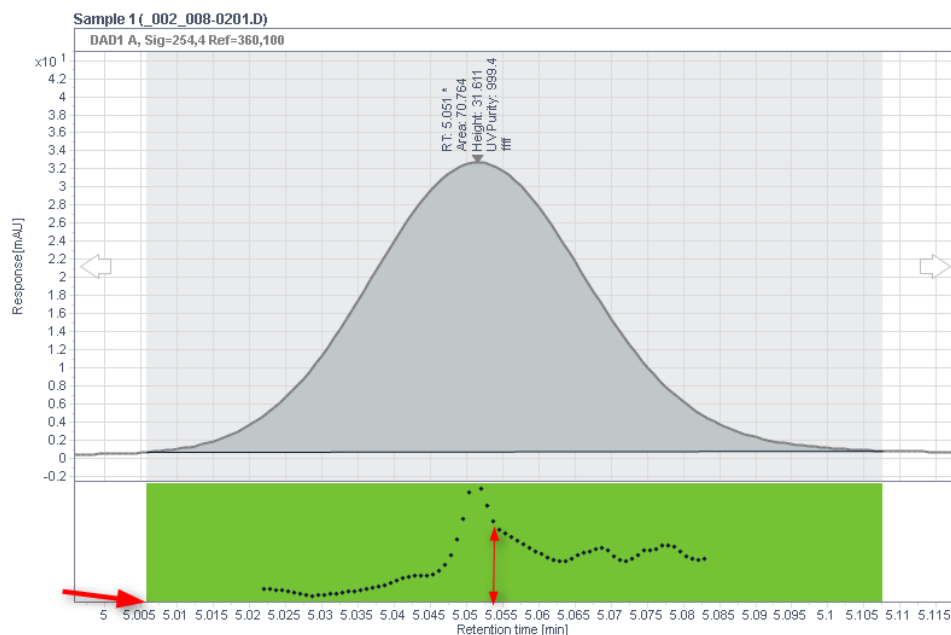
In OpenLab CDS ChemStation Edition, there are different threshold calculation types, in such way that thresholds can be fixed values or threshold curve types. Fixed threshold values do not consider the variation of the noise contribution across the peak. Therefore, in OpenLab CDS v2.x, this threshold is calculated automatically for every data point, leading to a threshold curve. The sharpness of the analysis is then adjusted using the sensitivity percentage.

The new plot shows the same information as in ChemStation, but in a more convenient manner. Instead of the classical threshold curve, it shows the logarithmic value of the difference (delta) of the threshold to the classical threshold curve. This creates a flat line, which is the border between the red and the green part of the plot. All points above the threshold are in the green area (black dots), and all points below the threshold are in the read area (white dots).

Compared to the ChemStation, the good and bad parts are flipped around: Good (green) is on the top, bad (red) is on the bottom. The bigger the distance from the threshold line in the green area, the better the value. The bigger the distance in the red area below the curve, the worse.

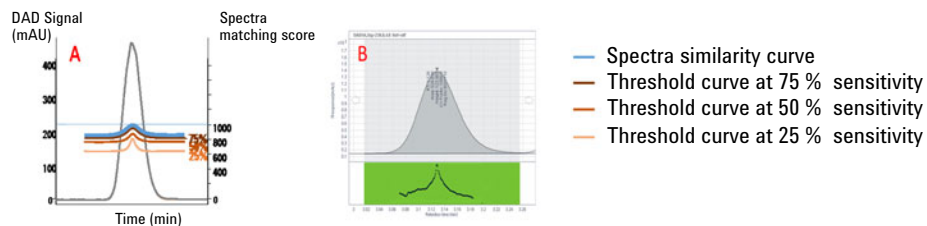
In cases where everything is above threshold, the threshold line is on the bottom of the plot:

Peak Details



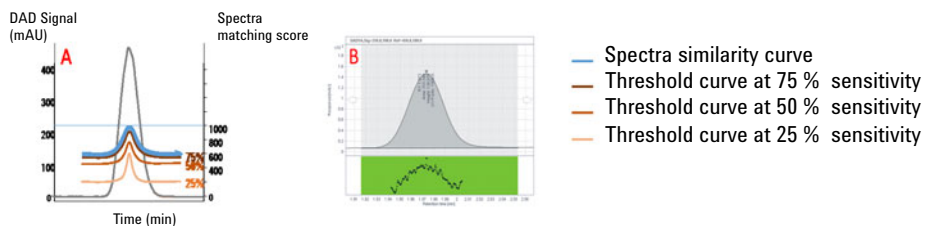
The following examples show the differences between a traditional purity plot, similar to the OpenLab CDS ChemStation edition similarity curve (which shows reversed similarity curves), and the OpenLab CDS v2.x similarity curve:

Pure peak



- | | |
|----|--|
| 1A | Schematic view of the similarity curve and its <i>calculated threshold curves</i> at different sensitivities |
| 2B | Data Analysis view of the similarity curve |

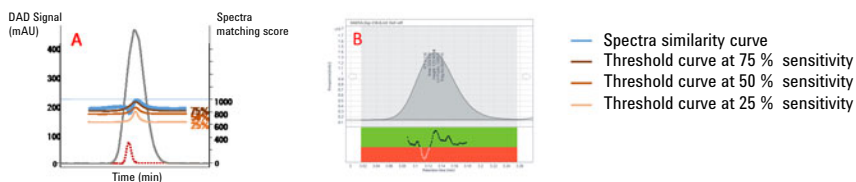
Pure peak with high noise value



1A Schematic view of the similarity curve and its *calculated threshold* curves at different sensitivities

2B Data Analysis view of the similarity curve

Impure peak



1A Schematic view of the similarity curve and its *calculated threshold* curves at different sensitivities

2B Data Analysis view of the similarity curve and its *logarithmic threshold line* at 50% sensitivity

Summary

Only threshold curve type methods can be compared between OpenLab CDS v2.x and OpenLab CDS ChemStation Edition, with the remaining differences that in OpenLab CDS v2.x reference background noise is automatically selected at peak start and peaks end, and the sensitivity is an additional factor which is not available in OpenLab CDS ChemStation Edition.

When using a default sensitivity of 50 %, the algorithm in both chromatographic data systems is the same, but the results will be different due to different noise references. If the sensitivity is increased or decreased, the threshold curve profiles will change compared to OpenLab CDS ChemStation Edition threshold curves.

UV confirmation

UV reference spectra are acquired from a reference sample under well-defined chromatographic conditions. You can confirm the identity of a compound by comparing the current spectrum at the peak apex with the UV reference spectrum. The application calculates a match factor for the two spectra.

The algorithm for the comparison is identical to the one used for the UV impurity check (see “[Similarity calculation](#)” on page 157). Background correction is optional, it can be selected in the processing method.

For UV confirmation, the match factor must be greater than a given limit to be colored green in the injection results. You set the match factor limit in the processing method. The resulting UV confirmation match factor is shown in the injection results.

Injection Results					
Peaks		Summary			
+	#	Name	Signal description	RT (min)	UV Conf. Match Factor
	1	A	DAD1A	0.895	1000
	2	B	DAD1A	1.330	1000
	3	C	DAD1A	1.789	1000
	4	D	DAD1A	2.272	693
	5	E	DAD1A	2.441	1000
	6	F	DAD1A	2.823	999

Figure 96 UV confirmation match factor in injection results



8 Mass Spectrometry

MS sample purity 168

MS peak purity 170

This chapter describes the sample purity calculation based on mass spectrometry.

MS sample purity

The MS sample purity calculation assesses whether a sample is pure or contains impurities. This assessment is based on the comparison of responses. On the one hand, there is the response of all compounds and fragments in a sample. On the other hand, there is the response caused by specific target ions. The sample purity is calculated as the ratio of both responses.

The application performs different steps to calculate the MS sample purity, depending on the selected base signal and calculation:

Target found?

- 1 Get the target masses given in the **Target** columns of the **Injection List** (for example, 270). If a formula is entered, calculate the molecular weight from the formula.
- 2 Apply the adducts specified in the processing method (for example, +H and +Na and a target mass of 270 would result in the targets 271 and 293).
- 3 Extract EICs for all targets, and sum these EICs to a single EIC.
- 4 Determine the retention time of the peak in that single summed EIC.
- 5 Locate the matching peak in the chromatogram of the base signal.

If a matching peak can be located, the target is marked as **found**.

Base signal is from an MS detector

With calculation **TIC %**

$$\text{MS sample purity} = \frac{\text{area of matching peak (TIC)}}{\text{area of all integrated peaks (TIC)}} \times 100$$

With calculation **EIC/TIC %**:

$$\text{MS sample purity} = \frac{\text{area of single peak (summed EIC)}}{\text{area of all integrated peaks (TIC)}} \times 100$$

NOTE

When **MS** is chosen as the base signal (in the processing method under **MS Sample Purity >Properties**), and there are multiple TIC signals, each TIC signal will be shown as a base signal on the left side table in the **Sample Purity Results** window.

Base signal is from another detector (non-MS)

$$\text{MS sample purity} = \frac{\text{area of matching peak (base signal)}}{\text{area of all integrated peaks (base signal)}} \times 100$$

Assumptions

MS sample purity is calculated under the following assumptions:

- The MS sample purity calculation is intended as a rough approximation only.
- For **EIC/TIC %** calculations: MS data is acquired such that most ion abundance is in the molecular ion cluster. There is only a small degree of in-source dissociation.
- For base signals from non-MS detectors: The other detector is more uniform and universal in its response than the MS detector.
- All compounds in the sample have uniform response factors.

MS peak purity

The MS peak purity calculation is based on the percentage of the component ions from spectra that group together and constitute the target relative to other components present.

The application performs the following steps to calculate the MS peak purity:

- 1 Run a deconvolution on the entire chromatographic range.
 - a For the top (n) detected m/z values, create an Extracted Ion Chromatogram (EIC). You can set the number in the processing method under **Compounds >Spectra, MS Peak Purity** tab.
 - b In each EIC, find the peak retention time.
 - c Define components based on EIC peaks that elute at the same retention time.
- 2 Determine the target component that matches the target compound.
 - a Get basic parameters for further calculations:
 - m/z delta range, from internal default settings (m/z -0.3 to +0.7)
 - Target quantifier m/z, by extracting the MS spectrum at the TIC peak apex and finding the m/z of highest abundance
 - Retention time window for the target compound (current compound), from the processing method under **Compounds >Identification**
 - b For each component found by deconvolution, find all m/z values that fall within the m/z delta range of the target quantifier m/z.
 - c For every such m/z value, check if the EIC has an apex within the retention time window of the target quantifier EIC peak.
 - d Get the component with the largest such peak, and use it as the target component.
- 3 If the system fails to find a target component, it re-tries by running the full-sample deconvolution in *high-resolution* mode, using an RT window size factor that is reduced by a factor of 2. The high-resolution results are cached and are searched for any target component that cannot be found in the normal-resolution component list. The automatic generation of the high-resolution results allows many target components to be identified that were previously missed.

- 4 The system attempts to detect *double components* that share the same RT and base peak m/z. The presence of such component doublets can strongly bias the purity estimate to the downside. Component doublets can occur if the RT window size factor is too small. Therefore, the system attempts to recover automatically by re-running the full-sample deconvolution in low-resolution mode, using a window size factor that is increased by a factor of 2. The low-resolution results are cached and are searched for any target that is matched to a component doublet.
- 5 Get all contributing components, that is, any component that has spectral peaks within m/z delta range of the target quantifier m/z and overlaps the target peak in retention time.
- 6 Calculate the purity using the following formula:

$$Purity(Target\ compound) = \frac{Area(Target\ component\ shape)}{\sum (Area(Contributing\ component\ shape))}$$

Evaluating system suitability	173
Noise Determination	175
Noise Calculation Using Six Times the Standard Deviation	176
Noise Calculation Using the Peak-to-Peak Formula	177
Noise Calculation by the ASTM Method	178
Noise calculation using the Root Mean Square (RMS)	181
Signal-to-noise calculation	182
Drift and Wander	184
Calculation of peak asymmetry and symmetry	186
System Suitability Formulas and Calculations	188
Performance Test Definitions	189
NOTE: Retention time used for performance test	189
Overview Performance Tests	190
True Peak Width W_x [min]	192
Capacity Factor (USP), Capacity Ratio (ASTM) k'	193
Symmetry Factor, Tailing Factor (USP) t	194
Number of Theoretical Plates per Column (USP)	195
Number of Theoretical Plates per Meter N [1-m]	195
Relative Retention, Selectivity	196
Resolution (USP, ASTM) R	196
Resolution (EP/JP) R_s	197
Peak to Valley Ratio (EP/JP)	197

This chapter describes what OpenLab CDS can do to evaluate the performance of both the analytical instrument and the analytical method.

Evaluating system suitability

Evaluating the performance of both the analytical instrument before it is used for sample analysis and the analytical method before it is used routinely is good analytical practice. It is also a good idea to check the performance of analysis systems before, and during, routine analysis. OpenLab provides the tools to do these three types of tests automatically. An *instrument test* can include the detector sensitivity, the precision of peak retention times and the precision of peak areas. A *method test* can include precision of retention times and amounts, the selectivity, and the robustness of the method to day-to-day variance in operation. A *system test* can include precision of amounts, resolution between two specific peaks and peak tailing.

Laboratories may have to comply with:

- Good Laboratory Practice regulations (GLP),
- Good Manufacturing Practice regulations (GMP) and Current Good Manufacturing Practice regulations (cGMP), and
- Good Automated Laboratory Practice (GALP).

Laboratories are advised to perform these tests and to document the results thoroughly. Laboratories which are part of a quality control system, for example, to comply with ISO9000 certification, will have to demonstrate the proper performance of their instruments.

To collate the results from several runs and evaluate them statistically, OpenLab CDS offers a function to create result set summary reports. Different report templates are available for these summaries (for example, SequenceSummary_Extended.rdl). They can be adjusted as required.

The tests are documented in a format which is generally accepted by regulatory authorities and independent auditors. Statistics include:

- peak retention time,
- peak area,
- amount,
- peak height,
- peak width at specific height,
- peak symmetry,
- peak tailing,
- capacity factor (k'),

- plate numbers,
- resolution between peaks, and
- selectivity relative to preceding peak.

Extended performance results are calculated only for calibrated compounds, ensuring characterization by retention times and compound names.

A typical system performance test report contains the following performance results:

- column details,
- processing method,
- sample information,
- acquisition information,
- signal description and baseline noise determination, and
- signal labeled with either retention times, or compound names.

In addition, the following information is generated for each calibrated compound in the chromatogram:

- retention/migration time,
- k' ,
- symmetry,
- peak width,
- plate number,
- resolution,
- signal-to-noise ratio, and
- compound name.

Noise Determination

Noise can be determined from the data point values from the time range of the current signal. Noise is treated in the following ways:

- as six times the standard deviation (sd) of the linear regression of the drift
- as peak-to-peak (drift corrected)
- as determined by the ASTM method (ASTM E 685-93)
- as the Root Mean Square (RMS) of the linear regression of the drift

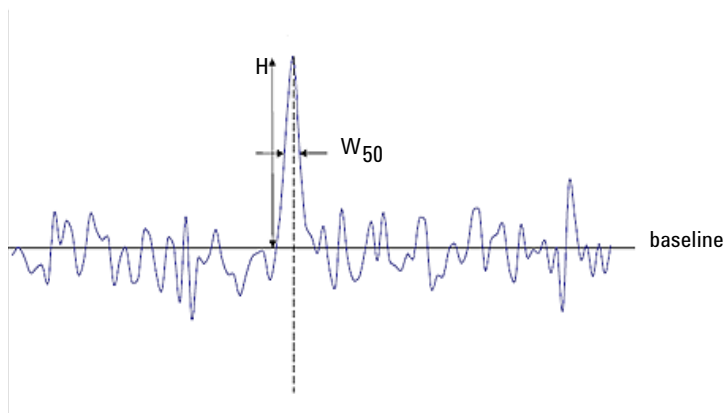


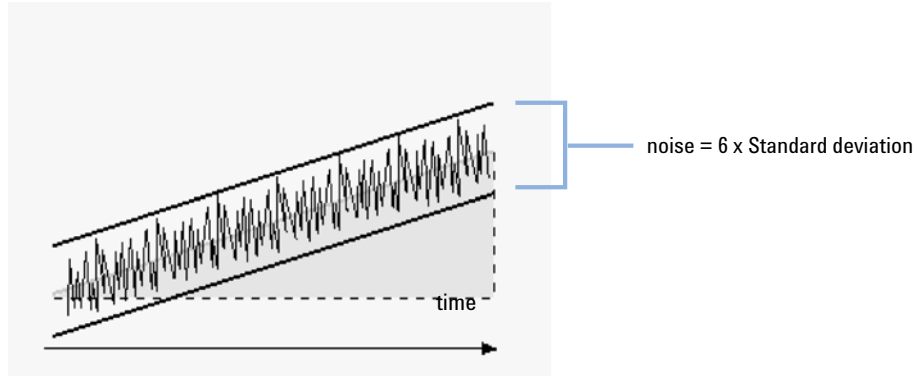
Figure 97 Chromatogram with peak signal and noise

H	Peak height from top to baseline (best straight line through noise)
W_{50}	Peak width at half height

NOTE

For very small peaks the application may find a retention time that is after peak end, which leads to a negative peak width. In this corner case, no noise value will be computed.

Noise Calculation Using Six Times the Standard Deviation



The linear regression is calculated using all the data points within the time range of the current signal. The noise is given by the formula:

$$N = 6 \times Std$$

where

N

Noise based on the six time standard deviation method

Std

Standard deviation of the linear regression of all data points in the selected time range

Noise Calculation Using the Peak-to-Peak Formula

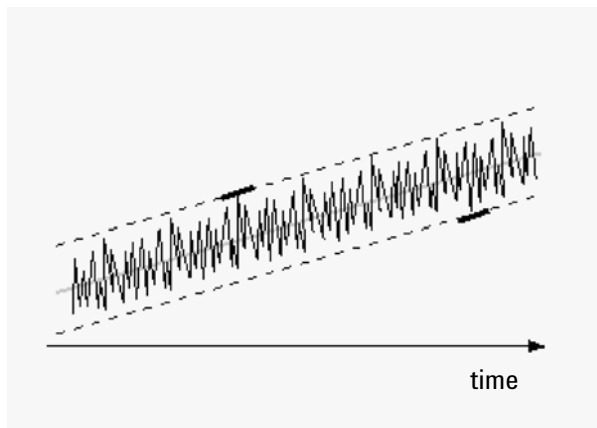


Figure 98 Illustration of peak-to-peak noise with drift

The drift is first calculated by determining the linear regression using all the data points in the time range of a peak. The linear regression line is subtracted from all data points within the time range to give the drift-corrected signal.

The peak-to-peak noise is then calculated using the formula:

$$N = I_{\max} - I_{\min}$$

where

N Peak-to-peak noise

$I_{\max}$ Highest (maximum) I_x value in the time range

$I_{\min}$ Lowest (minimum) I_x value in the time range

I_x Intensity of the signal, corrected by the drift (drift is calculated using the LSQ formula)

For European Pharmacopoeia calculations the Peak-to-Peak noise is calculated using a blank reference signal over a range of -10 and $+10$ times W_{50} flanking each peak. This region can be symmetrical to the signal of interest, or asymmetrical if required due to matrix signals.

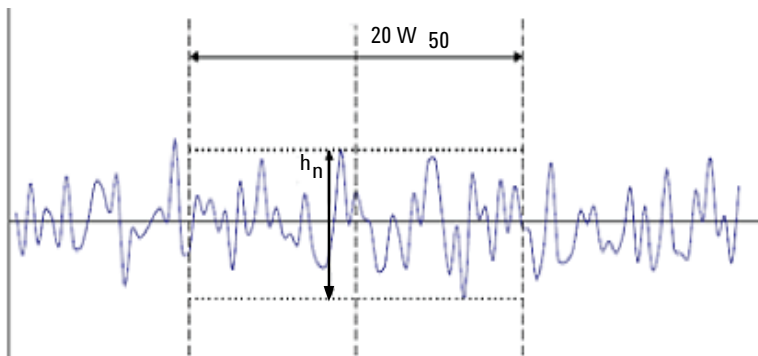


Figure 99 Determination of noise from the chromatogram of a blank sample

Where

$20 W_{50}$ is the region corresponding to the 20 fold of W_{50} .

h_n is the maximum amplitude of the baseline noise in the 20-fold W_{50} region.

Noise Calculation by the ASTM Method

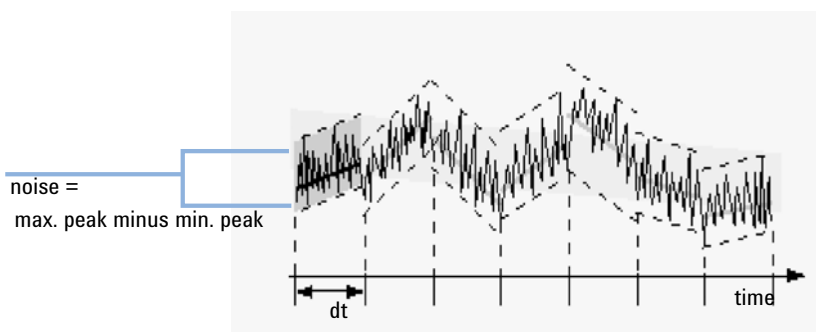


Figure 100 Noise determination by the ASTM method

ASTM noise determination (ASTM E 685-93) is based on the standard practice for testing variable-wavelength photometric detectors used in liquid chromatography, as defined by the American Society for Testing and Materials.

Based on the size of the time range, three different types of noise can be distinguished. Noise determination is based on peak-to-peak measurement within defined time ranges.

- *Cycle Time, t*

Long-term noise, the maximum amplitude for all random variations of the detector signal of frequencies between 6 and 60 cycles per hour. Long-term noise is determined when the selected time range exceeds one hour. The time range for each cycle (dt) is set to 10 minutes which will give at least six cycles within the selected time range.

Short-term noise, the maximum amplitude for all random variations of the detector signal of a frequency greater than one cycle per minute. Short-term noise is determined for a selected time range between 10 and 60 minutes. The time range for each cycle (dt) is set to one minute which will give at least 10 cycles within the selected time range.

Very-short-term noise (not part of ASTM E 685-93), this term is introduced to describe the maximum amplitude for all random variations of the detector signal of a frequency greater than one cycle per 0.1 minute.

Very-short-term noise is determined for a selected time range between 1 and 10 minutes. The time range for each cycle (dt) is set to 0.1 minute which will give at least 10 cycles within the selected time range.

- *Number of Cycles, n*

The number of cycles is calculated as:

$$n = \frac{t_{\text{tot}}}{t}$$

where t is the cycle time and t_{tot} is the total time over which the noise is calculated.

- *Peak-to-Peak Noise in Each Cycle*

The drift is first calculated by determining the linear regression using all the data points in the time range. The linear regression line is subtracted from all data points within the time range to give the drift-corrected signal. The peak-to-peak noise is then calculated using the formula:

$$N = I_{\max} - I_{\min}$$

where N is the peak-to-peak noise, $I_{\max}$ is the highest (maximum) intensity peak and $I_{\min}$ is the lowest (minimum) intensity peak in the time range.

- *ASTM Noise*

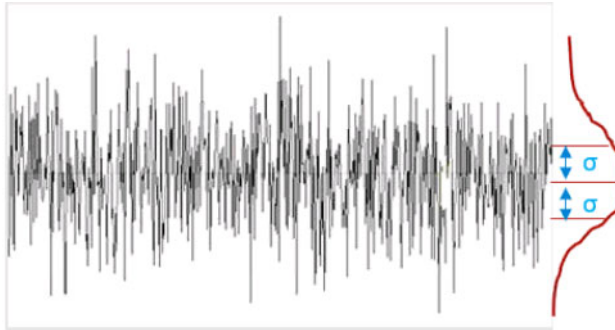
The ASTM noise is calculated as:

$$N_{\text{ASTM}} = \frac{\sum_{i=1}^n N}{n}$$

where N_{ASTM} is the noise based on the ASTM method.

An ASTM noise determination is not done if the selected time range is below one minute. Depending on the range, if the selected time range is greater than, or equal to one minute, noise is determined using one of the ASTM methods previously described. At least seven data points per cycle are used in the calculation. The cycles in the automated noise determination are overlapped by 10 %.

Noise calculation using the Root Mean Square (RMS)



The linear regression is calculated using all the data points within the time range of the current signal.

The noise is given by the formula:

$$\text{RMS} = S$$

where

RMS

Noise based on standard deviation method

S

Standard Deviation

Standard deviation of the linear regression of all data points in the selected time range, with linear function $y(X) = a + bX$:

$$S = \sqrt{\frac{\sum_{i=1}^N (Y_i - a - bX_i)^2}{N - 2}}$$

where

a

Y Intercept

b

slope

N

Number of discrete observations

X_i

Independent variable, i^{th} observation

Signal-to-noise calculation

OpenLab CDS has different options for the signal-to-noise calculation. You can choose both the algorithm and the noise range.

6 sigma or RMS method

The signal-to-noise is calculated using the formula:

$$\text{Signal-to-Noise} = \frac{\text{Height of the peak}}{\text{Noise of closest range}}$$

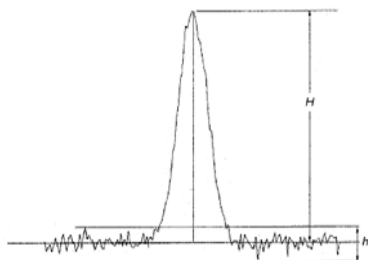


Figure 101 Signal-to-noise ratio

Peak to peak or ASTM method

The signal-to-noise is calculated using the formula:

$$S/N = 2H/h$$

where

H

Height of the peak corresponding to the component concerned in the chromatogram obtained with the prescribed reference solution.

h

Absolute value of the largest noise fluctuation from the baseline in a chromatogram obtained after injection or application of a blank and observed over a distance equal to twenty times the width at half-height of the peak in the chromatogram obtained with the prescribed reference solution, and situated equally around the place where this peak would be found.

According to the definition of the European Pharmacopoeia, signal-to-noise is calculated against a blank reference signal and a noise calculated over the time range which contains the peak the S/N ratio is being calculated for.

Noise range Noise can be calculated against the following time regions and signals:

- Fixed time region, on the same signal or on a blank reference signal
- Time region relative to the peak start or end, on the same signal or on a blank reference signal
- Automatically determined time region, on a blank reference signal.

An automatically determined time region is calculated according to one of the following algorithms:

- If the reference signal is not long enough
($EndTime - StartTime < 20 * W_{50}$)
 - $StartTime = starttime$ (of reference signal), and
 - $EndTime = endtime$ (of the reference signal)
- If the reference signal is long enough, but the peak is situated too close to the starttime
($t_R - 10 * W_{50} < starttime$ of the reference signal)
 - $StartTime = starttime$ (of reference signal), and
 - $EndTime = StartTime + 20 * W_{50}$
- If the reference signal is long enough, but the peak is situated too close to the endtime
($t_R + 10 * W_{50} > endtime$ of the reference signal)
 - $EndTime = endtime$ (of the reference signal), and
 - $StartTime = EndTime - 20 * W_{50}$
- If the reference signal is long enough, and the peak is situated far enough away from starttime and endtime of the reference signal
($t_R - 10 * W_{50} > starttime$, $t_R + 10 * W_{50} < endtime$)
 - $StartTime = t_R - 10 * W_{50}$, and
 - $EndTime = t_R + 10 * W_{50}$

where

t_R is the retention time, and

W_{50} is the peak width at half height.

Drift and Wander

Drift and wander are calculated if **Signal to noise** is selected in the processing method. They are calculated regardless of which noise calculation type you select.

Drift Drift is given as the slope of the linear regression. The drift is first calculated by determining the linear regression using all the data points in the time range. The linear regression line is subtracted from all data points within the time range to give the drift-corrected signal.

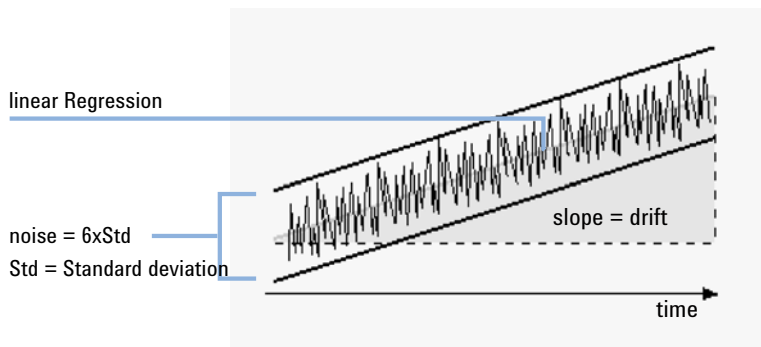


Figure 102 Drift for noise as Six Times the Standard Deviation

Curve Formula:

$$y(x) = a + bX$$

where

N	Number of discrete observations
X_i	Independent variable, i^{th} observation
Y_i	Dependent variable, i^{th} observation

Coefficients:

$$a = \frac{1}{\Delta X} \left(\sum_{i=1}^N X_i^2 * \sum_{i=1}^N Y_i - \left(\sum_{i=1}^N X_i * \sum_{i=1}^N X_i Y_i \right) \right)$$

$$b = \frac{1}{\Delta X} \left(N * \sum_{i=1}^N X_i Y_i - \left(\sum_{i=1}^N X_i * \sum_{i=1}^N Y_i \right) \right)$$

$$\Delta X = N * \sum_{i=1}^N X_i^2 - \left(\sum_{i=1}^N X_i \right)^2$$

Wander Wander is determined as the peak-to-peak noise of the mid-data values in the ASTM noise cycles, see [“Noise Calculation Using Six Times the Standard Deviation”](#) on page 176.

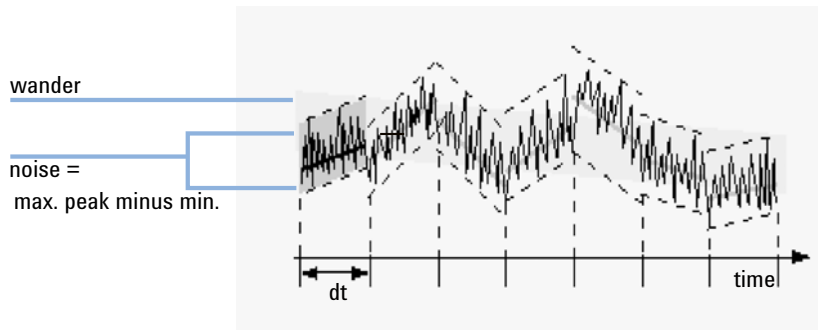


Figure 103 Wander of noise as determined by the ASTM Method

Calculation of peak asymmetry and symmetry

Asymmetry Peak asymmetry is calculated by comparing the peak half-widths at 10% of the peak height:

$$A_s = \frac{W_{10}}{2 W_{f, 10}}$$

where

A_s Asymmetry 10%

W_{10} Peak width at 10% of the peak height

$W_{f, 10}$ Front half of the peak width at 10% of the peak height.

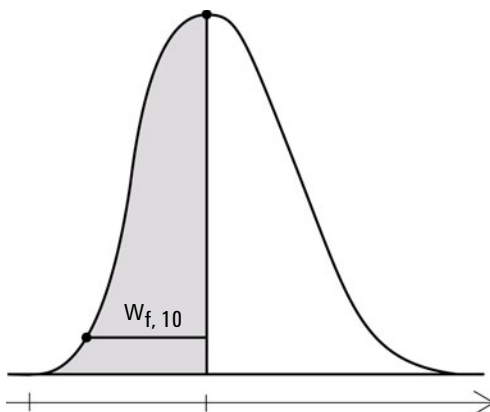


Figure 104 Calculation of peak asymmetry

Symmetry By most of the pharmacopeias, the symmetry factor of a peak is calculated by comparing the peak half-widths at 5%. In OpenLab, this factor is calculated and stored as the Tailing factor (see "[Symmetry Factor, Tailing Factor \(USP\) t](#)" on page 194). In OpenLab, the symmetry is calculated as a pseudomoment by the integrator using the following moment equations:

$$m_1 = a_1 \left(t_2 + \frac{a_1}{1.5 H_f} \right)$$

$$m_2 = \frac{a_2^2}{0.5 H_f + 1.5 H}$$

System Suitability

Calculation of peak asymmetry and symmetry

$$m_3 = \frac{a_3^2}{0.5H_r + 1.5H}$$

$$m_4 = a_4 \left(t_3 + \frac{a_4}{1.5H_r} \right)$$

$$\text{Peak symmetry} = \sqrt{\frac{m_1 + m_2}{m_3 + m_4}}$$

If no inflection points are found, or only one inflection point is reported, then the peak symmetry is calculated as follows:

$$\text{Peak symmetry} = \frac{a_1 + a_2}{a_3 + a_4}$$

where

a_i	Area of slice
t_i	Time of slice
H_f	Height of front inflection point
H_r	Height of rear inflection point
H	Height at apex

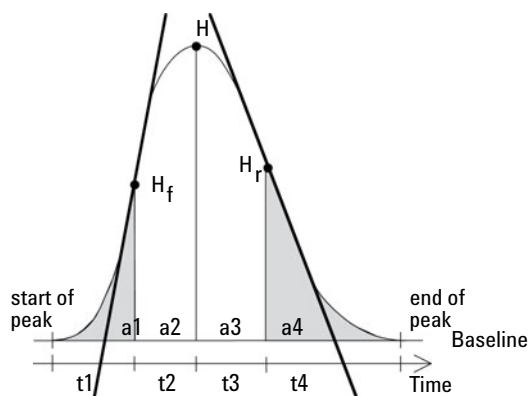


Figure 105 Calculation of the Peak Symmetry Factor

System Suitability Formulas and Calculations

The following formulas are used to obtain the results for the various System Suitability tests. The results are reported using the **Performance+Noise** and **Extended Performance** reports.

When ASTM or USP is specified for a given definition, then the definition conforms to those given in the corresponding reference. However, the symbols used here may not be the same as those used in the reference.

The references used in this context are:

- *ASTM: Section E 685 - 93, Annual Book of ASTM Standards, Vol. 14.01*
- *USP: The United States Pharmacopeia, XX. Revision, pp. 943 - 946*
- *EP: European Pharmacopoeia, 7th Edition*
- *JP: Japanese Pharmacopoeia, 16th Edition*

Performance Test Definitions

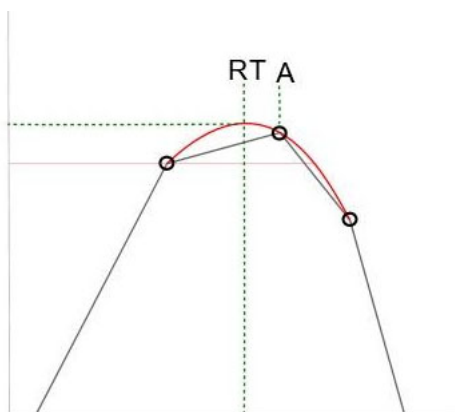
NOTE: Retention time used for performance test

Peak Performance can be calculated for any integrated peak of the data loaded, and also for new manually integrated peaks. The application calculates peak characteristics and column performance values using a peak model retention time which is calculated internally. It may differ slightly from the retention time shown in the injection results, chromatograms, or reports. The peak model retention time can be reported (see Data fields used for system suitability in the Reporting help, or search for *PeakModelIRT*).

NOTE

Please note that the retention time (RT) depicted in the figure below and determined by the peak integrator is not necessarily associated with the highest data point. The retention time is usually calculated using an parabolic interpolation model. This means that the retention time RT can potentially be smaller or larger than the retention time of the highest data point and the height might also be higher or lower than the highest data point.

The peak model retention time is calculated by smoothing the signal to avoid interference of noise. After this the highest data point closest to the RT of the peak is determined. The RT of this data point is used as the peak model retention time.



where

RT

Retention time shown in the injection results

A

Peak model retention time used for performance calculations

Overview Performance Tests

Table 9 Pharmacopeia values in OpenLab CDS

USP	EP	JP	Definition	Column name in injection results	Field used in reporting
Symmetry factor or tailing factor	Symmetry factor	Symmetry factor	$S = \frac{W_5}{2f}$	Tailing	Peak_TailFactor
Separation Factor	-	Separation Factor	$\alpha = \frac{k'_2}{k'_1} = \frac{t_{R2} - t_0}{t_{R1} - t_0}$	Selectivity	Peak_Selectivity
Relative retention	Relative retention	-	With RRT compounds: $\frac{RT_{peak} - t_0}{RT_{ref} - t_0}$	RRT EP	Peak_RelativeRetTime_EP
Relative retention time (RRT)	Unadjusted relative retention	Unadjusted relative retention	With RRT compounds $\frac{RT_{peak}}{RT_{ref}}$	RRT	Peak_RelativeRetTime
-	Resolution	Resolution	$R_s = 1.18 \cdot \frac{t_{R2} - t_{R1}}{W_{50(1)} + W_{50(2)}}$	Resol.EP Resol.JP	Peak_Resolution_EP Peak_Resolution_JP
Resolution	-	-	$R = 2 \cdot \frac{t_{R2} - t_{R1}}{W_{t(2)} + W_{t(1)}}$	Resol.USP	Peak_Resolution_USP
Number of theoretical plates (Efficiency)	-	-	$n = 16 \left(\frac{t_R}{W_t} \right)^2$	Plates USP	Peak_TheoreticalPlates_USP
-	Plate number (Efficiency)	Plate number (Efficiency)	$n = 5.54 \left(\frac{t_R}{W_{50}} \right)^2$	Plates EP Plates JP	Peak_TheoreticalPlates_EP Peak_TheoreticalPlates_JP

Table 9 Pharmacopeia values in OpenLab CDS

USP	EP	JP	Definition	Column name in injection results	Field used in reporting
S/N ratio	S/N ratio	S/N ratio	With P2P or ASTM noise calculation: $\frac{S}{N} = \frac{2H}{h}$ With 6SD or RMS noise calculation: $\frac{S}{N} = \frac{H}{h}$	S/N	Peak_SignalToNoise
Peak-to-valley ratio	Peak-to-valley ratio	Peak-to-valley ratio	$\frac{p}{v} = \frac{H_p}{H_v}$	p/v ratio	Peak_PeakValleyRatio

True Peak Width W_x [min]

W_x = width of peak at height x % of total

where

W_t	Tangent peak width, 4 sigma, obtained by intersecting tangents through the inflection points with the baseline
$W_{4.4}$	Width at 4.4% of height (5 sigma width)
W_5	Width at 5% of height (tailing peak width), used for USP tailing factor
W_{10}	Width at 10% of height
W_{50}	Width at 50% of height (true half-height peak width or 2.35 sigma).

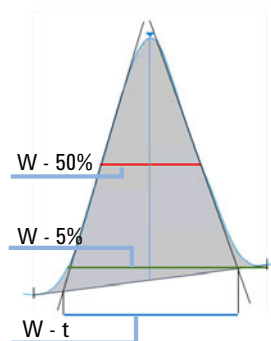


Figure 106 Peak width at x % of the height

Capacity Factor (USP), Capacity Ratio (ASTM) k'

$$k' = \frac{t_R - t_0}{t_0}$$

where

t_R

Retention time of peak [min]

t_0

Void time [min] (as provided in the processing method)

Symmetry Factor, Tailing Factor (USP) t

NOTE

The Symmetry Factor (USP, EP, JP) is identical with the Tailing Factor (USP). Both are available as "Peak_TailFactor" in Intelligent Reporting. See also [Table 9](#) on page 190.

$$S = \frac{W_5}{2f}$$

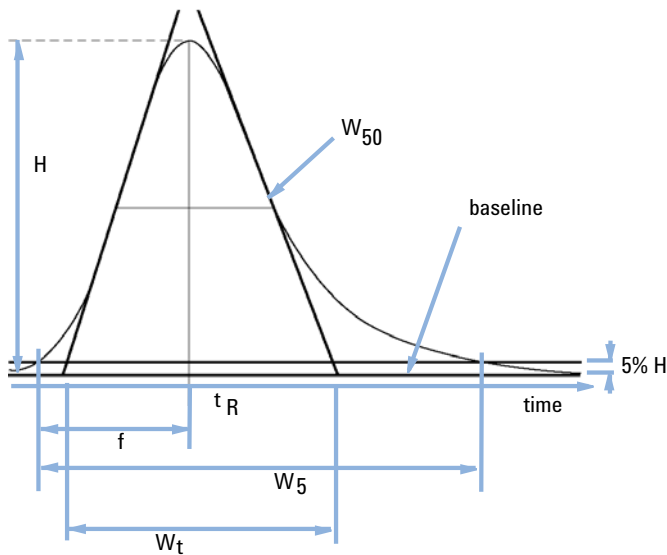


Figure 107 Performance Parameters

S	Symmetry factor, tailing factor (USP)
H	Peak height
t _R	Retention time
f	Distance in min between peak front and t _R , measured at 5% of the peak height
W ₅₀	Peak width at 50% of height [min]
W ₅	Peak width at 5% of peak height [min]
W _t	Tangent peak width

Number of Theoretical Plates per Column (USP)

Tangent method (USP, ASTM):

$$n = 16 \left(\frac{t_R}{W_t} \right)^2$$

where

t_R Retention time

W_t Tangent width [min]

Half-width method (ASTM, EP, JP):

$$n = 5.54 \left(\frac{t_R}{W_{50}} \right)^2$$

where

t_R Retention time

W_{50} Peak width at half-height [min]

Number of Theoretical Plates per Meter N [1-m]

$$N = 100 \cdot \frac{n}{l}$$

where

n Number of theoretical plates

l Length of column [cm] (as provided in the processing method)

Relative Retention, Selectivity

Selectivity

Selectivity calculates the alpha value for all signal peaks except the first one. For every pair of adjacent peaks (peaks 1 and 2, t_R of peak 1 < t_R of peak 2), the selectivity is calculated as follows:

$$\alpha = \frac{k'_{(2)}}{k'_{(1)}} = \frac{t_{R2} - t_0}{t_{R1} - t_0}, \alpha > 1$$

where

$k'_{(x)}$ capacity factor for peak x: $(t_{Rx} - t_0)/t_0$

Relative Retention (EP)

RRT (EP) can be calculated only if an RRT reference peak has been defined and identified. *Alpha* values are < 1 if the peak is to the left of the RRT reference, and > 1 if the peak is to the right of the RRT reference.

$$\frac{RT_{\text{peak}} - t_0}{RT_{\text{ref}} - t_0}$$

Resolution (USP, ASTM) R

Tangent method (pertaining to peaks 1 and 2, t_R of peak 1 < t_R of peak 2; t_R in min)

$$R = 2 \cdot \frac{t_{R2} - t_{R1}}{W_{t(2)} + W_{t(1)}}$$

where

t_R Retention time

W_t Tangent width [min]

Resolution (EP/JP) Rs

Resolution (JP) and Resolution (EP) are calculated with the half-width method (Resolution used in Performance Report):

$$R_s = 1.18 \cdot \frac{t_{R2} - t_{R1}}{W_{50(1)} + W_{50(2)}}$$

where

t_R Retention time

W_{50} Peak width at half-height [min]

Peak to Valley Ratio (EP/JP)

The peak to valley ratio (**p/v ratio** in the injection results) is calculated to indicate the quality of peak separation. It is calculated with the European and Japanese Pharmacopeia (EP, JP).

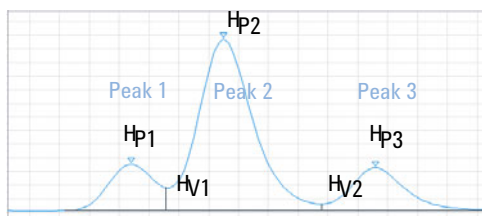
NOTE

This value is different from the peak to valley ratio that you set in the advanced integration parameters!

The peak to valley ratio is computed for peaks separated by a valley:

$PV = \text{peak height} / \text{valley height}$

If there are valleys to both left and right of a peak, the peak to valley ratio is computed for front and tail. The minimum p/v will be displayed.



For peak 1:

$$PV = \frac{H_{P1}}{H_{V1}}$$

For peak 2:

$$PV_F = \frac{H_{P2}}{H_{V1}}$$

$$PV_T = \frac{H_{P2}}{H_{V2}}$$

$$PV = \text{Min}(PV_F, PV_T)$$

For peak 3:

$$PV = \frac{H_{P3}}{H_{V2}}$$

where

PV	Peak to valley ratio
PV _F	Peak to valley ratio, front
PV _T	Peak to valley ratio, tail
H _{Px}	Height of peak x
H _{Vx}	Height of valley x

If a peak has multiple shoulders that are separated by valley, the peak to valley ratio is calculated for each shoulder.

Definition of a valley:

- Its height and time are shared between two consecutive peaks
- Its baseline is shared between two consecutive peaks
- The absolute baseline height is greater than 10⁻⁵.

The peak to valley calculation always uses absolute values. Therefore, even if one or more peaks are negative, the peak to valley ratio will always be shown as a positive number.

NOTE

The peak to valley ratio is not calculated for signals consisting of too few data points.

In This Book

This guide contains the reference information on the principles of operation, calculations and data analysis algorithms used in Agilent OpenLab CDS. The information contained herein may be used by validation professionals for planning and execution of system validation tasks.

- Integration with ChemStation algorithm
- Integration with EZChrom algorithm
- Peak Identification
- Calibration
- Quantitation
- UV Spectral Analysis
- Mass Spectrometry
- System Suitability