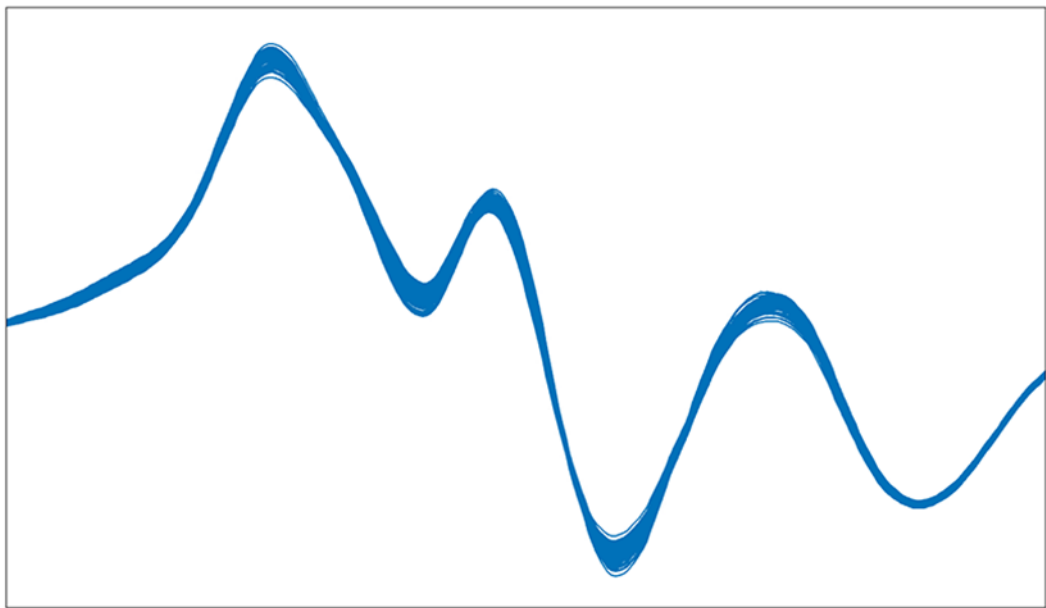


OMNIS NIR-Theorie



Handbuch

8.0600.8101DE / v10 / 2026-04-08



Metrohm AG
Ionenstrasse
CH-9100 Herisau
Schweiz
+41 71 353 85 85
info@metrohm.com
www.metrohm.com

OMNIS NIR-Theorie

Handbuch

8.0600.8101DE / v10 /
2026-04-08

Diese Dokumentation ist urheberrechtlich geschützt. Alle Rechte vorbehalten.

Bei dieser Dokumentation handelt es sich um ein Originaldokument.

Diese Dokumentation wurde mit grösster Sorgfalt erstellt. Dennoch sind Fehler nicht vollständig auszuschliessen. Bitte richten Sie diesbezügliche Hinweise an die obenstehende Adresse.

Haftungsausschluss

Von der Gewährleistung ausdrücklich ausgeschlossen sind Mängel, die auf Umstände zurückgehen, die nicht von Metrohm zu verantworten sind, wie unsachgemässe Lagerung, unsachgemässer Gebrauch etc. Eigenmächtige Veränderungen am Produkt (z. B. Umbauten oder Anbauten) schliessen jegliche Haftung des Herstellers für daraus resultierende Schäden und deren Folgen aus. Anleitungen und Hinweise in der Produktdokumentation der Metrohm sind strikt zu befolgen. Andernfalls ist die Haftung von Metrohm ausgeschlossen.

Inhaltsverzeichnis

1	Überblick	1
1.1	Einleitung	1
1.2	Konzeptrahmen	1
1.3	Angaben zur Dokumentation	2
1.4	Weiterführende Informationen	2
2	Nahinfrarotlicht und Spektren	3
2.1	Licht und seine Wechselwirkung mit Materie	3
2.2	Mathematische Grundlagen	6
2.2.1	Beer-Lambert-Gesetz	6
2.2.2	Lineare Regression	7
2.3	Wie Licht in ein Spektrum umgewandelt wird	9
3	Geräteeinrichtung	13
3.1	Wellenlängenkalibrierung	14
3.2	Referenzstandardisierung	15
3.2.1	OMNIS NIR Analyzer	16
3.2.2	2060 NIR	17
3.3	Geräteleistungstests	26
3.3.1	Externe Geräteleistungstests (OMNIS NIR Analyzer)	30
4	Modellentwicklung	32
4.1	Physische Proben	34
4.2	Hauptkomponentenanalyse (PCA)	36
4.3	Datenaufbereitung	41
4.3.1	Datenvorbehandlung	41
4.3.2	Wellenlängenbereiche	50
4.3.3	Spektrale Ausreisser	52
4.3.4	Referenzwert-Ausreisser (Quantifizierung)	59
4.3.5	Datensatzaufteilung	60
4.4	Quantifizierung	62
4.4.1	PLS-Regression	62
4.4.2	Validierung von Quantifizierungsmodellen	65
4.4.3	OMNIS Model Developer (OMD)	73
4.4.4	Steigungs-/y-Achsenabschnittskorrektur	74
4.5	Identifizierung und Verifizierung	77
4.5.1	Support Vector Machine (SVM)	77
4.5.2	Prädiktion der Produktzugehörigkeit einer Probe	80
4.5.3	Validierung von Identifizierungsmodellen	82

1 Überblick

1.1 Einleitung

Die Nahinfrarotspektroskopie (NIR-Spektroskopie) ist eine zerstörungsfreie, schnelle und reagenzienfreie Analysenmethode, die sich für ein breites Spektrum von Proben eignet. Sie kann mehrere Parameter zugleich analysieren und sowohl physikalische als auch chemische Eigenschaften eines Materials bestimmen. Dazu gehören unter anderem Analytkonzentrationen, Dichte, Partikelgrösse oder intrinsische Viskosität.

Die NIR-Spektroskopie ermöglicht ausserdem die Identifizierung unbekannter Proben (ab OMNIS Software 4.0) und die Verifizierung von Proben (ab OMNIS Software 4.2).

Die Möglichkeit zur zerstörungsfreien Messung von Proben aus der Ferne ist in der Qualitätskontrolle und Prozessüberwachung von entscheidender Bedeutung.

Das Handbuch beschreibt Techniken und Algorithmen zur Aufnahme, Verarbeitung und Analyse von Nahinfrarot-Spektren, wie sie in der OMNIS Software implementiert sind. In Kapitel 2 wird kurz erläutert, wie die Messsignale in Absorptionsspektren umgewandelt werden. Kapitel 3 behandelt die Kalibrierung des Geräts. Kapitel 4 beschreibt die Entwicklung von Modellen, die den Analyseparameter (Quantifizierung) oder die Produktzugehörigkeit (Identifizierung) vorhersagen können. Kapitel 5 befasst sich mit der Prädiktion von Unbekannten. Kapitel 6 bildet den Anhang mit Erläuterungen zu verschiedenen Algorithmen.

1.2 Konzeptrahmen

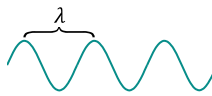
Die vorgestellten Verfahren fügen sich in den folgenden Rahmen ein:

1. **Kalibrierung, Standardisierung und Leistungstests**
Sicherstellung der Übertragbarkeit und Zuverlässigkeit der mit dem Gerät erfassten Absorptionsspektren.
2. **Modellentwicklung**
Es wird ein Modell für die Prädiktion eines quantitativen Parameters oder für die Identifizierung von Proben entwickelt.
Die Entwicklung basiert auf Proben mit bekanntem Analyseparameter oder bekannter Produktzugehörigkeit.

2 Nahinfrarotlicht und Spektren

2.1 Licht und seine Wechselwirkung mit Materie

Ein Spektrometer misst, wie eine Probe mit Licht interagiert. Licht kann in unterschiedlichem Masse absorbiert oder gestreut werden. Die Wechselwirkung hängt von den Eigenschaften des Lichts, insbesondere seiner Wellenlänge, und von den Eigenschaften des Materials, insbesondere seiner Molekularstruktur, ab.



Wellenlänge

Licht ist elektromagnetische Strahlung. Das Licht bewegt sich als Welle mit schwingenden elektrischen und magnetischen Feldern durch den Raum. Die Wellen breiten sich in Raum und Zeit aus. Eine Welle wird dementsprechend durch ihre Wellenlänge λ (z. B. in Nanometer = 10^{-9} Meter) und ihre Frequenz (Hz) charakterisiert.

Die Wellenlänge ist umgekehrt proportional zur Frequenz der Welle. Wellen mit höheren Frequenzen (mehr Schwingungen pro Sekunde) haben kürzere Wellenlängen und umgekehrt. Aufgrund dieser Beziehung kann die Welle entweder mit der Wellenlänge (nm) oder der Frequenz (Hz) beschrieben werden.

Licht kann Energie in diskreten Quanteneinheiten, den Photonen, austauschen. Die Energie eines einzelnen Photons, E , hängt von seiner Frequenz f bzw. seiner Wellenlänge λ ab:

$$E = hf = h \frac{c}{\lambda}$$

Hierbei entspricht h der Planck-Konstante und c der Lichtgeschwindigkeit.

Abbildung 1 zeigt verschiedene Bereiche von elektromagnetischer Strahlung.

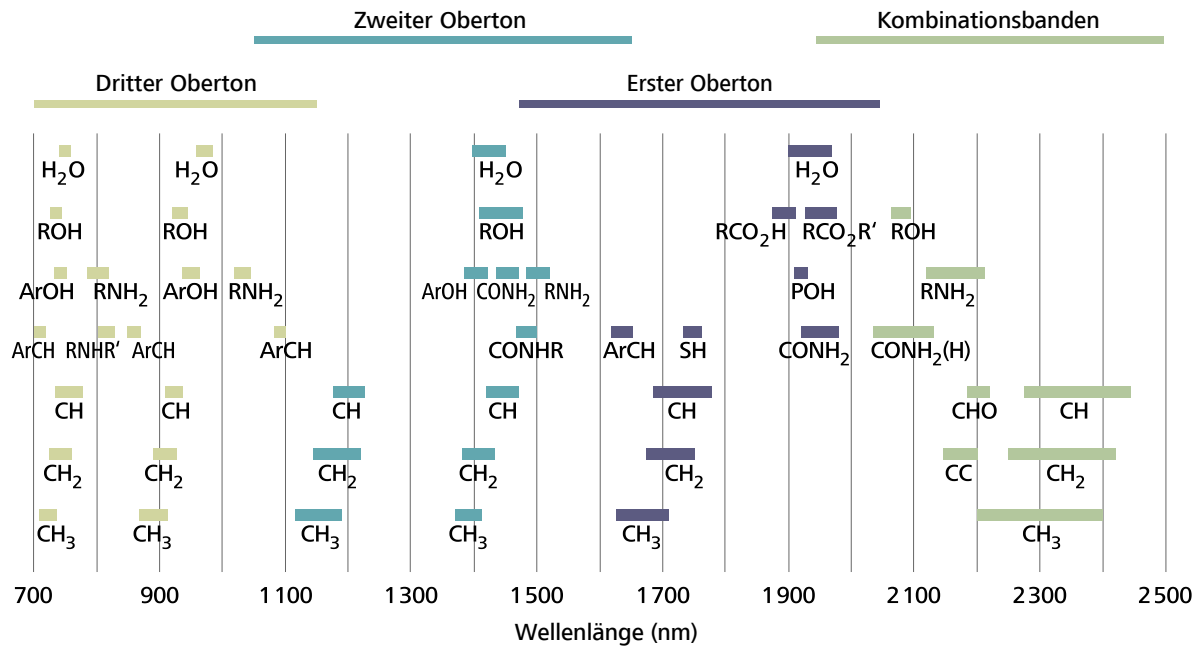


Abbildung 2 NIR-Absorptionsbanden

Der Grundübergang ist der wahrscheinlichste Übergang, er kommt am häufigsten vor. Obertonübergänge sind weniger wahrscheinlich. Daher absorbiert der Grundübergang mehr Licht als Obertonübergänge. Im Allgemeinen nimmt die Absorption mit jedem Oberton ab. Dadurch sind die Obertöne für stark absorbierende Moleküle geeignet.

Zwei oder mehr Grundschwingungen können gleichzeitig mit einer einzigen Lichtfrequenz angeregt werden, die den kombinierten Frequenzen der Grundschwingungen entspricht. Die entsprechenden Absorptionsbanden werden als **Kombinationsbanden** bezeichnet. Einige Kombinationsbanden liegen im NIR-Bereich, nämlich zwischen 1'900 und 2'500 nm.

2.2 Mathematische Grundlagen

2.2.1 Beer-Lambert-Gesetz

Das Beer-Lambert-Gesetz beschreibt, wie die Absorption von Licht durch eine homogene Probe von den Eigenschaften einer absorbierenden Substanz in der Probe abhängt.

$$A = \varepsilon \cdot c \cdot l$$

Hierbei entspricht A der Absorbanz, ε dem molaren Extinktionskoeffizient des Absorbers (L/mol/cm), c der Konzentration des Absorbers (mol/L) und l der Schichtdicke der Probe (cm).

Der molare Extinktionskoeffizient ε ist eine Konstante, die angibt, wie viel eine Substanz absorbiert. Der molare Extinktionskoeffizient ist spezifisch für eine bestimmte Wellenlänge i und eine bestimmte Substanz j . Die

GesamtabSORBANZ eines Gemisches ist die Summe der AbsORBANZ aller im Gemisch enthaltenen Substanzen:

$$A_i = \sum_{j=1}^N \varepsilon_{ij} c_j l$$

Hierbei entspricht A_i der AbsORBANZ bei der Wellenlänge i , N der Anzahl der Substanzen im Gemisch, ε_{ij} dem molaren Extinktionskoeffizienten für die Wellenlänge i und die Substanz j und c_j der Konzentration der Substanz j .

Das Beer-Lambert-Gesetz setzt eine lineare Beziehung zwischen der AbsORBANZ und der Konzentration sowie eine lineare Beziehung zwischen der AbsORBANZ und dem molaren Extinktionskoeffizienten voraus. Dieses lineare Verhalten gilt für viele Situationen.

Auf Grundlage des Beer-Lambert-Gesetzes können spektroskopische AbsORPTIONSMESSUNGEN Folgendes nachweisen:

- Schwankungen in der Konzentration eines AbsORBERS.
Das ist die häufigste Anwendung.
- Schwankungen bei den Faktoren, die den molaren Extinktionskoeffizienten beeinflussen.
Temperatur, Viskosität, pH-Wert oder die Dielektrizitätskonstante des Lösungsmittels können den molaren Extinktionskoeffizienten beeinflussen. In einigen Fällen kann das für spektroskopische Messungen genutzt werden.

Nicht mit dem Beer-Lambert-Gesetz verbunden sind Streueffekte. Streueffekte können manchmal genutzt werden, um z. B. Schwankungen bei der Partikelgrösse zu erkennen.

2.2.2 Lineare Regression

Eine Wellenlänge

Im einfachsten Fall hat ein Gemisch nur einen AbsORBER. Nach dem Beer-Lambert-Gesetz ist die AbsORBANZ bei einer bestimmten Wellenlänge linear zur Konzentration des AbsORBERS.

In Abbildung 1 stellt jeder Punkt eine Probe mit einer bekannten Konzentration des AbsORBERS (x-Achse) und der gemessenen AbsORBANZ (y-Achse) dar. Eine lineare Regression ergibt die Regressionsgerade **A**.

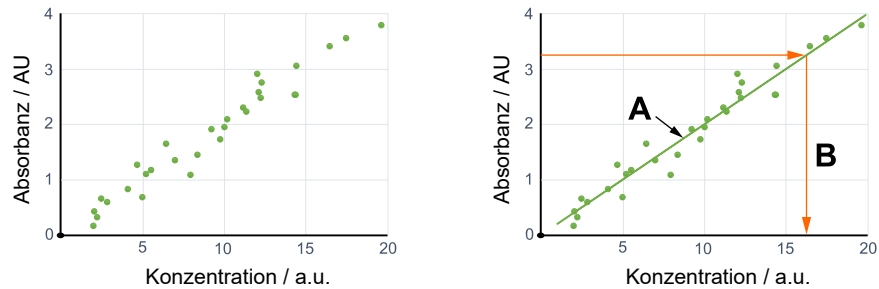


Abbildung 3 Beziehung zwischen Absorbanzwerten und Konzentrationen

Für eine Probe mit unbekannter Konzentration des Absorbers kann die Konzentration wie folgt bestimmt werden:

1. Messung der Absorbanz bei der vorgegebenen Wellenlänge.
2. Verwendung der Regressionsgeraden zur Ermittlung der Konzentration (**B**).

Die Regressionsgerade ist ein einfaches Quantifizierungsmodell zur Prädiktion des Analyseparameters, z. B. der Konzentration des Absorbers.

Diese Vorgehensweise schlägt jedoch fehl, wenn das Gemisch mehrere Absorber mit unterschiedlichen Konzentrationen enthält.

Mehrere Wellenlängen

Anstelle einer Messung der Absorption bei einer einzigen Wellenlänge kann die Absorption bei mehreren Wellenlängen gemessen werden, was dann ein Spektrum ergibt. Ähnlich wie oben kann mittels einer linearen Regression die Beziehung zwischen den Spektren und dem Analyseparameter ermittelt werden. Für mehrere Wellenlängen ist eine multiple lineare Regression (MLR) erforderlich.

Es ist nachweisbar, dass die multiple lineare Regression den Analyseparameter auch dann vorhersagen kann, wenn das Gemisch mehrere Absorber mit unterschiedlichen Konzentrationen enthält (*siehe "Beispiel für eine lineare Regression", Kapitel 6.1, Seite 91*).

Für eine multiple lineare Regression wäre jedoch eine grössere Probenanzahl erforderlich als die Anzahl der Wellenlängen. Für spektroskopische Prädiktionen können andere Methoden wie PCA oder PLS verwendet werden.

2.3 Wie Licht in ein Spektrum umgewandelt wird

Ein Spektrometer (oder Spektrophotometer) besteht aus einer Lichtquelle und einer Detektoreinheit. Die Lichtquelle strahlt Licht mit einem breiten Spektrum an Wellenlängen ab, mit anderen Worten polychromatisches Licht. Das Licht interagiert mit der Probe. Anschliessend erfasst das Spektrometer das verbleibende Licht als eine Funktion der Wellenlänge.

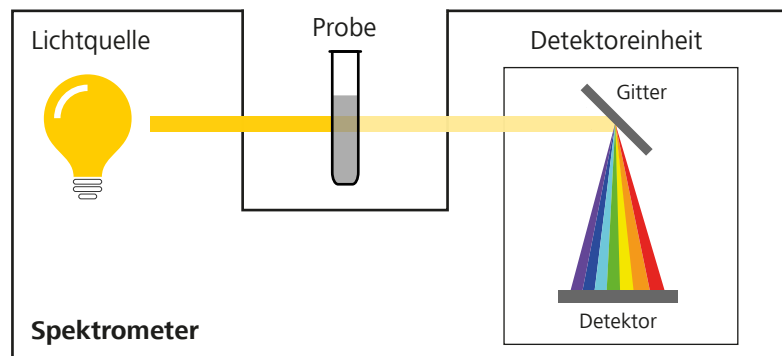


Abbildung 4 Ein Spektrometer mit Lichtquelle und Detektoreinheit.

Im Spektrometer wird das Licht mit Hilfe eines Gitters in einzelne Wellenlängen zerlegt. Ein Detektor misst das Licht, wobei jede Wellenlänge auf ein anderes Element – oder Pixel – in einem Zeilensensor trifft.

Ein **Scan** ist eine Messung über alle Pixel. Jedes Pixel erzeugt ein fotoelektrisches Signal, das proportional zur Lichtintensität ist. Die Signale können gegen die Pixel aufgetragen werden.

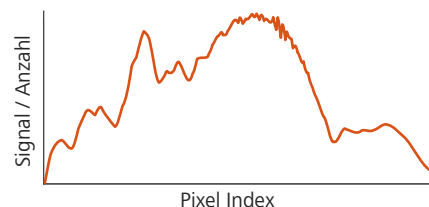


Abbildung 5 Spektrum des Detektorsignals als Funktion der Pixel.

Integrationszeit

Die Integrationszeit ist die Zeitspanne, in welcher der Detektor Licht sammelt. Eine höhere Integrationszeit erhöht das Signal.

Zu lange Integrationszeiten führen zu einer Sättigung des Detektors und zu einem Verlust an Information. Zu kurze Integrationszeiten verringern das Signal und damit das Signal-Rausch-Verhältnis.

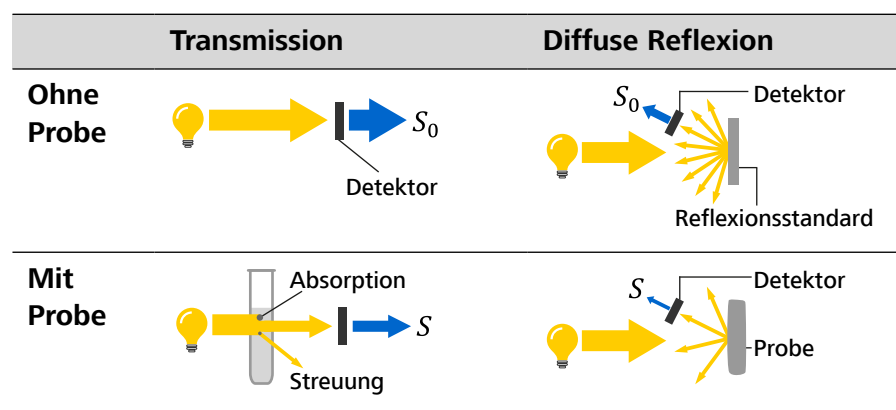
Eine **automatische Integrationszeit** sorgt für eine optimale Belichtung, d. h. ein optimales Signal-Rausch-Verhältnis, ohne dass es zu einer Sättigung kommt. Vor jedem Probenscan und jedem Referenzscan werden mehrere Messungen durchgeführt. Die Integrationszeit wird so eingestellt,

$$A = \log_{10} \frac{I_0}{I} = \log_{10} \frac{S_0}{S}$$

Transmission und Reflexion

Im **Transmissionsmodus** wird das durch die Probe hindurchtretende Licht gemessen. S_0 wird in Abwesenheit der Probe gemessen. S wird aus dem Licht gemessen, das durch die Probe gedrungen ist.

Im **Reflexionsmodus** wird das von der Probe reflektierte Licht gemessen. Als Referenz wird anstelle der Probe ein Reflexionsstandard verwendet. Der Reflexionsstandard reflektiert idealerweise 100 % des Lichts. Ein Teil des reflektierten Lichts wird zum Detektor geleitet und liefert das Signal S_0 . Das Signal S wird auf die gleiche Weise gemessen, jedoch mit der Probe, die das Licht reflektiert.



Die berechnete Absorbanz A repräsentiert alles Licht, das den Detektor nicht erreicht. Daher beinhaltet A nicht nur das von der Probe absorbierte Licht, sondern auch:

- Licht, das den Detektor nicht erreicht, weil es vom Detektor weg gestreut wird.
- Licht, das fälschlicherweise zum Detektor gestreut wird.

Absorptionsspektrum

Das Absorptionsspektrum wird anhand des Referenzscans (Signal S_0) und des Probenscans (Signal S) nach der obigen Formel berechnet.

3 Geräteeinrichtung

Die folgenden Schritte stellen sicher, dass für eine bestimmte Probe mit einem bestimmten Probenmessaufbau identische Spektren gewonnen werden, unabhängig davon, ob sie zu unterschiedlichen Zeiten und mit demselben oder einem anderen Gerät aus der Produktfamilie **OMNIS NIR Analyzer** oder einem anderen Gerät des Typs **2060 NIR** aufgenommen wurden.

Sowohl die x-Achse als auch die y-Achse der Spektren müssen berücksichtigt werden:

- **x-Achse:** Wellenlängenkalibrierung (*siehe "Wellenlängenkalibrierung", Kapitel 3.1, Seite 14*)
- **y-Achse:** Referenzstandardisierung (*siehe "Referenzstandardisierung", Kapitel 3.2, Seite 15*)

Darüber hinaus muss sichergestellt werden, dass die Geräteleistung den Anforderungen entspricht, d. h.:

- Die **Geräteleistungstests** müssen erfolgreich durchgeführt werden, bevor mit dem Gerät Spektren aufgenommen werden können (*siehe "Geräteleistungstests", Kapitel 3.3, Seite 26*).

Standards

Für die Wellenlängenkalibrierung und die Geräteleistungstests nutzt das Gerät einen internen, metrologisch rückführbaren **Wellenlängenstandard**.

Bei der Nutzung des Reflexionsmodus wird zudem ein **Reflexionsstandard** für die Referenzstandardisierung und die Geräteleistungstests benötigt, der sich nach dem Gerätetyp richtet:

- **OMNIS NIR Analyzer:** interner Reflexionsstandard
- **2060 NIR:** externer Reflexionsstandard

3. Validierung der Bandbreiten:
 - a. Im aufgenommenen Spektrum werden die Peakbreiten bestimmt.
 - b. Die Bandbreitenresiduen zwischen den gemessenen Peakbreiten und den bekannten Peakbreiten werden berechnet.
 - c. Für jeden Peak muss das Bandbreitenresiduum innerhalb der Toleranz liegen, um den Test zu bestehen.
4. Der Gesamtstatus der Validierung ist erfolgreich, wenn alle oben genannten Residuen innerhalb der Toleranz liegen.

Die Validierung muss erfolgreich durchgeführt werden, bevor mit dem Gerät Spektren aufgenommen werden können .

3.2 Referenzstandardisierung

Die Referenzstandardisierung normiert die Absorbanzwerte, d. h. die y-Achse der Spektren.

Bestimmung der Absorbanz

Für die Berechnung der Absorbanz A einer Probe werden die Signale S_0 (Referenzscan) und S (Probenscan) benötigt (siehe "Wie Licht in ein Spektrum umgewandelt wird", Kapitel 2.3, Seite 9):

$$A = \log_{10} \frac{S_0}{S}$$

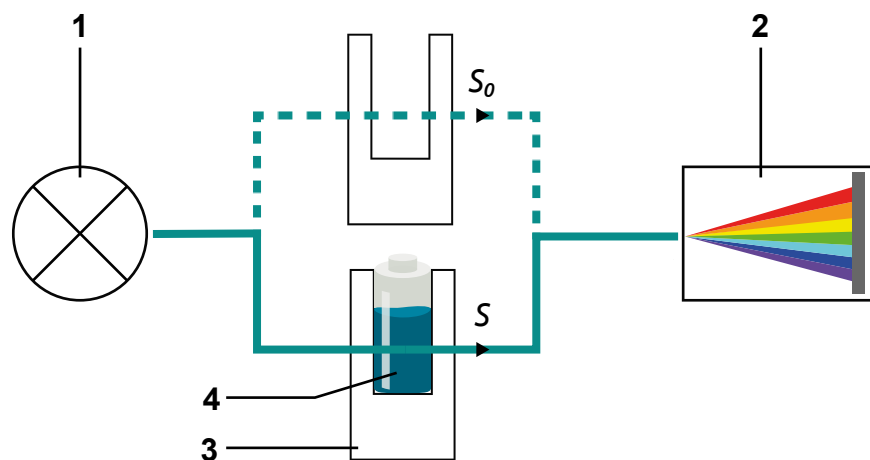


Abbildung 8 Optischer Pfad im Transmissionsmodus (als Beispiel mit einer Flüssig-Probenpräsentation).

In [Abbildung 8](#) gelangt das Licht von der Lichtquelle (1) durch einen Probenhalter (3) zum Detektor (2).

Das Referenzsignal S_0 wird ohne die Probe gemessen, das Signal S mit der Probe (4). Ansonsten sind die optischen Eigenschaften für beide optischen

3.2.2 2060 NIR

Geräte vom Typ **2060 NIR** erfordern eine externe Referenzstandardisierung.

Externe Referenzstandardisierung

Die wiederholte Messung des Signals S (mit der Probe) und des Signals S_0 (ohne die Probe) über optische Pfade mit identischen optischen Eigenschaften ist zeitaufwändig und fehleranfällig.

Deshalb werden zwei weitere Strahlengänge eingeführt (*siehe Abbildung 9, Seite 17*):

- Eine **interne Referenz** im Gerät. Der interne Referenzpfad liefert das Signal S_{ref} , das auf einfache Weise gemessen werden kann.
- Ein weiterer externer optischer Pfad, bei dem die Faseroptiken mit einer **Kalibriervorrichtung** verbunden sind. Dieser optische Pfad liefert das Signal S_{fiber} .

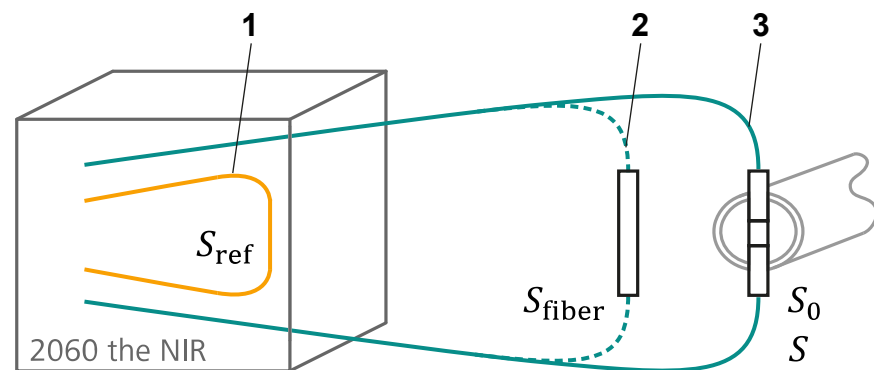


Abbildung 9 Optische Pfade als Beispiel im Transmissionsmodus: Interner Referenzpfad (1), externe Faseroptiken verbunden mit einer Kalibriervorrichtung (2) und externe Faseroptiken verbunden mit der Sonde mit oder ohne Probe (3). Die optischen Pfade 2 und 3 stellen die gleiche Faseroptik dar, die lediglich anders verbunden ist.

Die Kalibriervorrichtung fixiert die Faseroptiken und bildet damit den Referenzpfad (2). Im Transmissionsmodus dient die Luft als Referenz und überträgt 100 % des Lichts. Im Reflexionsmodus nimmt die Kalibriervorrichtung auch den Reflexionsstandard auf. Zunächst wird von einem idealen Reflexionsstandard ausgegangen, der 100 % des Lichts reflektiert.

Die Absorbanz A der Probe wird aus den Signalen S_0 und S berechnet. Werden die beiden zusätzlichen Signale S_{ref} und S_{fiber} zum Zähler und Nenner hinzugefügt, bleibt das Resultat unverändert:

$$A = \log_{10} \frac{S_0}{S} = \log_{10} \left(\frac{S_{\text{ref}}}{S} \cdot \frac{S_{\text{fiber}}}{S_{\text{ref}}} \cdot \frac{S_0}{S_{\text{fiber}}} \right)$$

Validieren der Glasfaserkorrektur

Die Glasfaserkorrektur muss mit denselben Messparametern und derselben Kalibriervorrichtung validiert werden.

Bei diesem Verfahren wird ein Referenzmaterial verwendet:

- Im Reflexionsmodus ist das Referenzmaterial ein Reflexionsstandard. Es wird von einem nicht idealen Reflexionsstandard ausgegangen (z. B. 99 %). Der Reflexionsstandard hat ein bekanntes nominales Absorptionsspektrum A_{nominal} .
- Im Transmissionsmodus dient die Luft als Referenz. Das nominale Absorptionsspektrum ist eine Nulllinie ($A_{\text{nominal}} = 0$), da angenommen wird, dass die Luft kein Licht absorbiert.

Abbildung 12 veranschaulicht, wie die Residuen der Validierung bestimmt werden:

1. Die externen Faseroptiken müssen an die Kalibriervorrichtung angeschlossen werden.
Im Reflexionsmodus wird die Kalibriervorrichtung mit dem Reflexionsstandard kombiniert.
2. Der Befehl **VAL REF STD** mit der Schnittstelle **Glasfaser** führt die folgenden Scans durch:
 - a. Ein interner Referenzscan liefert einen Wert für S_{ref} .
 - b. Ein externer Scan auf dem jeweiligen Kanal misst die externen Faseroptiken, die Kalibriervorrichtung und das Referenzmaterial. Dies ergibt das Signal S_{raw} .
3. Die Software berechnet A_{raw} (**Gemessenes Rohspektrum**):

$$A_{\text{raw}} = \log_{10} \frac{S_{\text{ref}}}{S_{\text{raw}}}$$
4. A_{raw} wird durch das Glasfaserkorrekturspektrum korrigiert, um die Absorbanz der Faseroptiken und der Kalibriervorrichtung zu eliminieren:

$$A_{\text{corrected}} = A_{\text{raw}} - A_{\text{fiber}}$$
 $A_{\text{corrected}}$ wird in der Software als **Gemessenes korrigiertes Spektrum** angezeigt.
5. Im Idealfall sollte $A_{\text{corrected}}$ mit dem **Referenzspektrum** A_{nominal} übereinstimmen. Die Differenzen zwischen den beiden werden als **Residuen der Validierung** berechnet:

$$A_{\text{residual}} = A_{\text{corrected}} - A_{\text{nominal}}$$

Hinweis: Im Transmissionsmodus ist $A_{\text{nominal}} = 0$ und $A_{\text{residual}} = A_{\text{corrected}}$.

1. Um einen aktuellen Wert für A_{fiber} zu erhalten, sollte eine Glasfaserkorrektur wie oben beschrieben durchgeführt werden.
Wichtig: Die Glasfaserkorrektur muss vor der Fensterkorrektur erfolgen.
2. Die externen Faseroptiken müssen dann ohne vorhandene Probe an die Sonde angeschlossen werden. Bei Bedarf nimmt ein Reflexionsstandard den Platz der Probe ein.
3. Der Befehl **REF STD** mit der Schnittstelle **Fenster** führt die folgenden Scans durch:
 - a. Ein interner Referenzscan liefert einen Wert für S_{ref} .
 - b. Ein externer Scan auf dem jeweiligen Kanal misst die externen Faseroptiken, die Sonde und das Referenzmaterial. Dies ergibt das Signal S_{probe} .
4. Die Absorbanz A_{probe} bezogen auf den internen optischen Pfad ist:

$$A_{\text{probe}} = \log_{10} \frac{S_{\text{ref}}}{S_{\text{probe}}}$$
5. Die Software berechnet A_{raw} (**Gemessenes Rohspektrum**):

$$A_{\text{raw}} = A_{\text{probe}} - A_{\text{fiber}}$$

Hierbei entspricht A_{raw} der Absorbanz der Sonde und des Referenzmaterials, bezogen auf die Kalibriervorrichtung.
6. Das nominale Absorptionsspektrum des Referenzmaterials, A_{nominal} , wird in der Software als **Referenzspektrum** angezeigt. Das Referenzspektrum muss von A_{raw} subtrahiert werden, um A_{window} zu erhalten:

$$A_{\text{window}} = A_{\text{raw}} - A_{\text{nominal}}$$

Hierbei entspricht A_{window} der Absorbanz der Sonde, bezogen auf die Kalibriervorrichtung.

Hinweis: Im Transmissionsmodus ist $A_{\text{nominal}} = 0$ und $A_{\text{window}} = A_{\text{raw}}$.

A_{window} steht für das Fenster**Korrekturspektrum**.
7. A_{window} bleibt unverändert, bis wieder eine Fensterkorrektur für den jeweiligen Kanal durchgeführt wird.

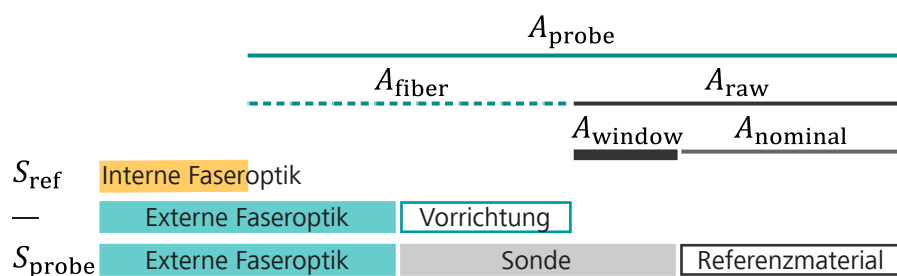


Abbildung 13 Fensterkorrektur

Validieren der Fensterkorrektur

Die Fensterkorrektur muss mit denselben Messparametern und derselben Kalibriervorrichtung validiert werden.

- **2060 NIR**

Die Tests können den internen optischen Pfad oder einen der externen optischen Pfade nutzen. Die Tests prüfen die Wellenlängen und das Signalrauschen.

Vor den Tests müssen die Wellenlängenkalibrierung und die externe Referenzstandardisierung auf dem jeweiligen Kanal erfolgreich durchgeführt und validiert werden.

Die Geräteleistungstests müssen erfolgreich durchgeführt werden, bevor mit dem Gerät Spektren auf dem jeweiligen Kanal aufgenommen werden können. Vordefinierte Toleranzen müssen eingehalten werden. Die zulässigen Toleranzen hängen von der Faseroptikkonfiguration des jeweiligen Kanals ab, die in den gerätespezifischen Daten angegeben ist (Messmodus, Fasertyp und Faserlänge).

Wellenlängentest

Beim Wellenlängentest werden Wellenlängengenauigkeit und Wellenlängenpräzision untersucht. Dazu wird ein Wellenlängenstandard verwendet, der ein Absorptionsspektrum mit definierten Peaks und bekannten Peakpositionen aufweist:

- **Intern:** Das Absorptionsspektrum des internen, metrologisch rückführbaren Wellenlängenstandards wird über den internen optischen Pfad mehrfach bestimmt:

$$A_{WL,x} = \log_{10} \left(\frac{S_{\text{ref}}}{S_{\text{ref},WL,x}} \right)$$

Hierbei entspricht $A_{WL,x}$ dem Absorptionsspektrum des internen Wellenlängenstandards für die x-te Messung, S_{ref} dem auf dem internen Referenzpfad gemessenen Referenzspektrum und $S_{\text{ref},WL,x}$ dem auf dem internen Referenzpfad mit eingesetztem internen Wellenlängenstandard gemessenen Spektrum der x-ten Messung.

- **Extern** für Geräte der Produktfamilie **OMNIS NIR Analyzer**: Der externe Wellenlängentest ist optional (*siehe "Externer Wellenlängentest", Seite 30*).

Rauschtest

Das Signalrauschen kann intern oder extern geprüft werden:

- **Intern:** Das Rauschen wird mehrfach über die Absorbanz des internen optischen Pfads bestimmt, bezogen auf die Absorbanz der Referenzmessung auf demselben optischen Pfad:

$$A_{\text{noise},x} = \log_{10} \left(\frac{S_{\text{ref}}}{S_{\text{ref},x}} \right)$$

Hierbei entspricht $A_{\text{noise},x}$ dem Rauschspektrum der x-ten Messung, S_{ref} dem auf dem internen Referenzpfad gemessenen Referenzspektrum und $S_{\text{ref},x}$ dem auf demselben Pfad gemessenen Spektrum der x-ten Messung.

- **Extern** für Geräte der Produktfamilie **OMNIS NIR Analyzer**: Der externe Rauschtest ist optional (*siehe "Externe Rauschtests", Seite 30*).
- **Extern** für Geräte vom Typ **2060 NIR**:
Die externen Faseroptiken müssen bei Transmission an die Kalibriervorrichtung und bei Reflexion an den Reflexionsstandard angeschlossen sein.

Das Rauschen wird als Differenz zwischen dem gemessenen Absorptionsspektrum und dem nominalen Absorptionsspektrum mehrfach bestimmt:

$$A_{\text{noise},x} = \log_{10} \left(\frac{S_{\text{ref}}}{S_{\text{fiber},x}} \right) - A_{\text{fiber}} - A_{\text{nominal}}$$

Hierbei entspricht $A_{\text{noise},x}$ dem Rauschspektrum der x-ten Messung, S_{ref} dem auf dem internen Referenzpfad gemessenen Referenzspektrum, $S_{\text{fiber},x}$ dem auf dem optischen Pfad des jeweiligen Kanals gemessenen Spektrum der x-ten Messung, wobei die Fasern mit der Kalibriervorrichtung bzw. dem Reflexionsstandard verbunden sind, A_{fiber} dem Glasfaserkorrekturspektrum aus der Referenzstandardisierung und A_{nominal} dem nominalen Absorptionsspektrum des Reflexionsstandards.

Hinweis: Im Transmissionsmodus ist $A_{\text{nominal}} = 0$.

Im Idealfall ist $A_{\text{noise}} = 0$.

Der Rauschtest führt die folgenden Schritte durch:

1. Das Spektrum $S_{\text{ref},x}$ bzw. $S_{\text{fiber},x}$ wird 10 mal gemessen, woraus sich 10 Rauschspektren $A_{\text{noise},x}$ ergeben.
2. Die Rauschspektren werden in verschiedene Wellenlängensegmente unterteilt.
3. Für jedes Rauschspektrum und jedes Segment werden 3 Größen berechnet:
 - a. Photometrisches Rauschen (Einheit: mAU)
 - b. Peak-To-Peak-Rauschen (Einheit: mAU)
 - c. Basislinien-Bias des Rauschens (Einheit: mAU)

4 Modellentwicklung

Folgende Arten von Modellen werden unterschieden:

- Ein **Quantifizierungsmodell** beschreibt die Abhängigkeit eines Analyseparameters (z. B. Wassergehalt) von den aufgenommenen Spektren der Proben.
- Ein **Identifizierungsmodell** (ab OMNIS Software-Version 4.0) klassifiziert die Proben anhand der aufgenommenen Spektren in verschiedene Produkte (z. B. verschiedene Sorten von Kaffeebohnen). Ein Produkt bezeichnet eine bestimmte chemische Substanz oder eine bestimmte chemische Substanz mit spezifischen physikalischen Eigenschaften (z. B. Partikelgrösse).

Für die Analyse einer unbekannten Probe wird ein Spektrum der Probe aufgenommen. Je nach Anwendung wird das Spektrum wie folgt genutzt:

- Quantifizierung: Auf der Grundlage des Spektrums erstellt ein Quantifizierungsmodell eine Prädiktion, z. B. für den Wassergehalt der Probe.
- Identifizierung: Auf der Grundlage des Spektrums identifiziert ein Identifizierungsmodell die Probe, z. B. als Arabica-Kaffee.
- Verifizierung (ab OMNIS Software-Version 4.2): Auf der Grundlage des Spektrums verifiziert ein Identifizierungsmodell beispielsweise, ob es sich bei der Probe um Arabica-Kaffee handelt oder nicht.

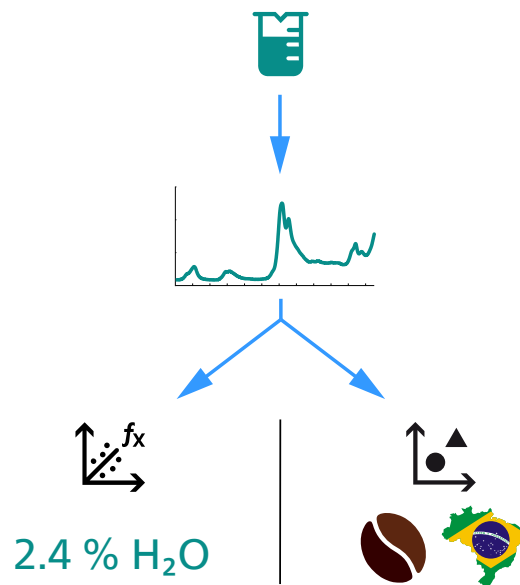


Abbildung 16 Quantifizierung (unten links) und Identifizierung (unten rechts).

Die Entwicklung eines Modells und die Analyse von Proben umfasst die folgenden Schritte:

Probenanzahl

Je mehr Variationen bei den Bedingungen, chemischen Komponenten oder Partikelgrößen abgedeckt werden müssen, umso mehr Proben werden benötigt.

Quantifizierung: Damit die statistische Analyse ordnungsgemäss funktioniert, ist eine Mindestanzahl von etwa 50 Proben erforderlich, wobei Kalibrierset und Validierset jeweils mindestens 20 bis 25 Proben enthalten müssen.

Identifizierung und Verifizierung: Für jedes Produkt müssen die Proben die erwarteten Variationen abdecken. Die Produkte können eine unterschiedliche Probenanzahl haben, die Mindestanzahl beträgt 3.

Mit mindestens 10 bis 20 Proben (je nach Anzahl der Variationen) kann ein erstes Modell ohne Validierset entwickelt werden. Falls die Kreuzvalidierung (für Quantifizierungsmodelle) oder die interne Validierung (für Identifizierungsmodelle) darauf hindeutet, dass ein adäquates Modell erstellt werden kann, müssen für die Entwicklung des endgültigen Modells weitere Kalibrier- und Validierproben gesammelt werden.

Replikate

Manchmal gibt es nur sehr wenige Proben in einem bestimmten Bereich von Bedingungen oder Referenzwerten. Um das auszugleichen, könnte man versucht sein, diese Proben zu replizieren. Das ist jedoch problematisch. Wenn sich Duplikate einer Probe sowohl im Kalibrierset als auch im Validierset befinden, sind die statistischen Kennzahlen irreführend (zu optimistisch). Auch Duplikate im gleichen Set sind zu vermeiden.

Referenzmethode (Quantifizierung)

Bei der Quantifizierung wird eine Referenzmethode zur Messung der Referenzwerte verwendet. Der **Standardfehler des Labors (SEL)** für die verwendete Referenzmethode spielt bei der Entwicklung eines Quantifizierungsmodells eine bedeutende Rolle. Der SEL ist die Standardabweichung der Differenzen zwischen den Messungen von Duplikatproben.

Der SEL ist häufig der grösste Fehler, der zum Standardfehler der Prädiktion (SEP) bei der NIR-Methode beiträgt (*siehe "SEP – Standardfehler der Prädiktion", Seite 72*). Der SEL sollte das 0.7-fache, vorzugsweise das 0.5-fache, des erforderlichen SEP nicht überschreiten. Der Bereich der Referenzwerte sollte mindestens das 3-fache, besser das 5-fache des SEL betragen.

Eine Möglichkeit zur Reduzierung des SEL ist die Durchführung wiederholter Referenzmessungen an jeder Probe. Der Durchschnitt der Messwerte sollte als Referenzwert für die Probe festgelegt werden. Für jede Probe sollte die gleiche Anzahl an Referenzmessungen durchgeführt werden. Die statistischen Kennzahlen werden relativ zu einer bestimmten Anzahl wie-

i Die OMNIS Software nutzt die PCA während der Modellentwicklung für die automatische Datensatzaufteilung und die Erkennung spektraler Ausreisser.

Vorbereitende Schritte

Folgende vorbereitende Schritte sind erforderlich:

1. **Datenvorbehandlung:** Die OMNIS Software wendet die festgelegte Datenvorbehandlung auf die Spektren an (*siehe "Datenvorbehandlung", Kapitel 4.3.1, Seite 41*).
2. **Wellenlängenbereiche:** Die OMNIS Software wendet die festgelegte Wellenlängenauswahl auf die Spektren an (*siehe "Wellenlängenbereiche", Kapitel 4.3.2, Seite 50*).
3. **Mittelwertzentrierung:** Für jede Wellenlänge wird der mittlere Absorbanzwert berechnet und vom jeweiligen Wert in jedem Spektrum subtrahiert.

Die erste Hauptkomponente

Nach den vorbereitenden Schritten ordnet die PCA die Informationen in den Spektrendaten neu an und trennt die relevanten Daten vom Rauschen. Zu diesem Zweck transformiert die PCA die Wellenlängenvariablen in einen neuen Variablenraum, in sogenannte **Hauptkomponenten (PC, englisch Principal Components)**.

Die PCA wandelt die relevanten Informationen aus einer Vielzahl von Wellenlängenvariablen in nur wenige Hauptkomponenten um. Um das Konzept anhand eines einfachen Beispiels zu veranschaulichen, wird angenommen, dass es statt Tausender nur 2 Wellenlängenvariablen gibt und diese 2 Variablen auf 1 Hauptkomponente reduziert werden.

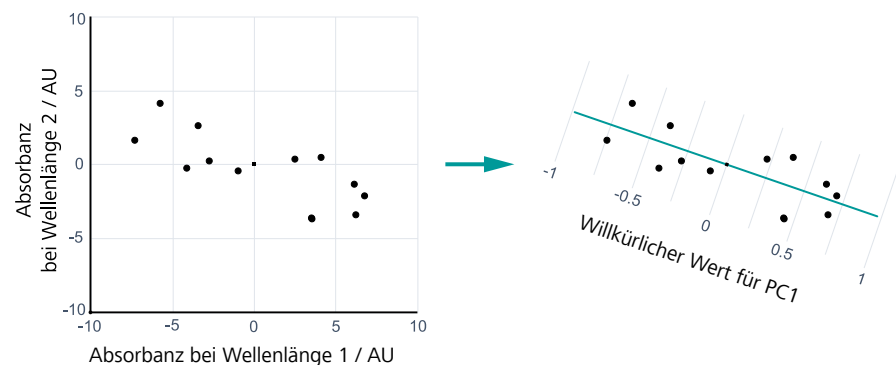


Abbildung 17 Punkte, die Spektren in einem 2-dimensionalen Wellenlängenraum darstellen (links). Die gleichen Punkte in einem 1-dimensionalen Hauptkomponentenraum (rechts).

In *Abbildung 17* links bilden die horizontale und die vertikale Achse den ursprünglichen Wellenlängenraum mit 2 Variablen. Somit stellt jeder Punkt

- Ein wellenlängenabhängiger multiplikativer Faktor zweiter Ordnung, der zu einer **quadratischen Neigung der Basislinie** führt. Es können auch Basislinienneigungen höherer Ordnung auftreten.
- Absorbanzabhängige multiplikative Faktoren, die zu einer **Verstärkung** führen. Verstärkungen sind jedoch irrelevant.

Die Streuung ist bei festen Proben am offensichtlichsten. Die sich daraus ergebenden Basislinienverschiebungen können verwendet werden, um Änderungen der Partikelgröße oder andere physikalische Variationen zu erkennen.

Sind jedoch chemische Variationen von Interesse, sollten die Basislinienverschiebungen durch eine geeignete Vorbehandlung minimiert werden.

Datenvorbehandlungen

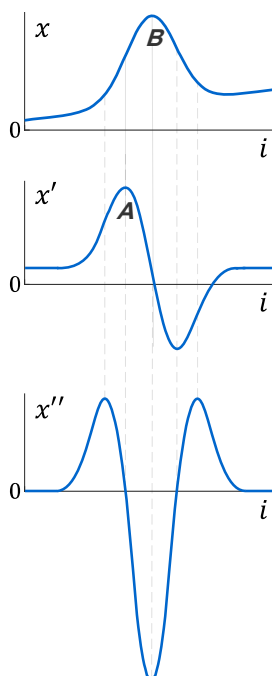
i Alle nachfolgenden Datenvorbehandlungen werden jeweils auf ein einzelnes Spektrum angewendet. Es werden keine weiteren Spektren in die Berechnungen einbezogen.

Die Datenvorbehandlung verändert die Signalwerte. Die Zahlen auf der y-Achse haben keine Bedeutung mehr.

Bei den angewandten Datenvorbehandlungen handelt es sich um lineare Transformationen. Daher gilt weiterhin das Beer-Lambert-Gesetz.

4.3.1.1 Ableitungen

i Ableitungen können mit dem Gap-Segment-Filter oder dem Savitzky-Golay-Filter durchgeführt werden.



Die Ableitung eines Spektrums beschreibt die Steigung oder Steilheit der Kurve an jedem Punkt. Die Steigung ist die Änderungsrate des Ausgangsspektrums.

Im Spektrum ist x_i die Absorbanz bei der Wellenlänge i . Die Ableitung erster Ordnung x'_i stellt die Steigung des Spektrums bei der Wellenlänge i dar. Wo das Ausgangsspektrum am steilsten ist, hat die Ableitung erster Ordnung ein Maximum (**A**). Wo das Ausgangsspektrum einen Peak (**B**) aufweist, ist die Ableitung erster Ordnung gleich 0.

Die Ableitung erster Ordnung beseitigt Offsets der Basislinie und wandelt Neigungen der Basislinie in Offsets der Basislinie um.

Die Ableitung zweiter Ordnung x''_i entspricht der Steigung der Ableitung erster Ordnung bei der Wellenlänge i . Positive Peaks im Originalspektrum (**B**) werden zu negativen Peaks und umgekehrt.

Die Ableitung zweiter Ordnung beseitigt Offsets der Basislinie und Neigungen der Basislinie aus dem Originalspektrum.

Vorsicht ist geboten, wenn das Originalspektrum ein erhebliches Mass an Rauschen enthält. Jede Ableitung verschlechtert das Signal-Rausch-

Parameter

Wellenlängenbereiche

Falls Artefakte gewisse Wellenlängenbereiche beeinträchtigen (z. B. durch Sättigung oder starkes Rauschen), können diese Bereiche ausgeschlossen werden.

Für die Berechnung des Mittelwerts und der Standardabweichung werden ausschliesslich die definierten Wellenlängenbereiche berücksichtigt. Die anschliessende Normierung wird sowohl in den definierten Wellenlängenbereichen als auch in den Zwischenbereichen durchgeführt. Für ausgeschlossene Wellenlängen am Anfang des Spektrums wird der normierte Wert der benachbarten Startwellenlänge übernommen. Für ausgeschlossene Wellenlängen am Ende des Spektrums wird der normierte Wert der benachbarten Endwellenlänge übernommen. Bei Bedarf können ausgeschlossene Wellenlängen auch für die Berechnung des Modells ausgeschlossen werden (*siehe "Wellenlängenbereiche", Kapitel 4.3.2, Seite 50*).

Hinweis: Ab OMNIS Software-Version 4.6 können mehrere Wellenlängenbereiche definiert werden.

4.3.1.5 Detrend

Detrend passt mithilfe der Methode der kleinsten Quadrate ein Polynom zweiter Ordnung an das gesamte Spektrum an. Anschliessend subtrahiert Detrend das Polynom vom Spektrum.

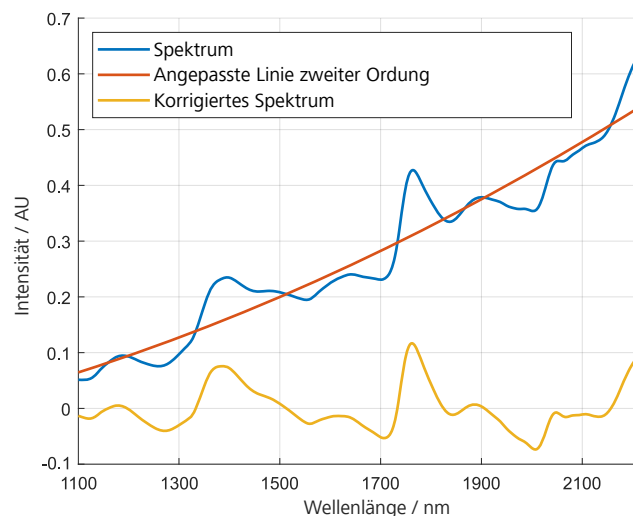


Abbildung 25 Detrend wandelt das blaue Spektrum in das gelbe Spektrum um.

Detrend reduziert wellenlängenabhängige Streueffekte bis hin zu quadratischen Neigungen der Basislinie.

Die obige Abbildung zeigt ein Spektrum (blau), in dem der Trend dominiert. Falls dieser dominierende Trend bei allen Spektren ähnlich ist, kann

- **OMNIS NIR Analyzer**

Die Integrationszeit wird immer automatisch eingestellt. Dies verhindert Sättigung und minimiert das Rauschen (*siehe "Integrationszeit", Seite 9*).

- **2060 NIR**

Falls die automatische Integrationszeit aktiviert ist, tritt keine Sättigung auf.

Ist die manuelle Integrationszeit aktiviert, können zu lange Integrationszeiten zur Sättigung des Detektors führen. Gesättigte Bereiche treten bei niedrigen Absorbanzwerten auf, sind aber visuell nicht immer leicht zu erkennen. Die manuelle Integrationszeit sollte daher mit ausreichend Spielraum eingestellt werden (*siehe "Integrationszeit", Seite 9*).

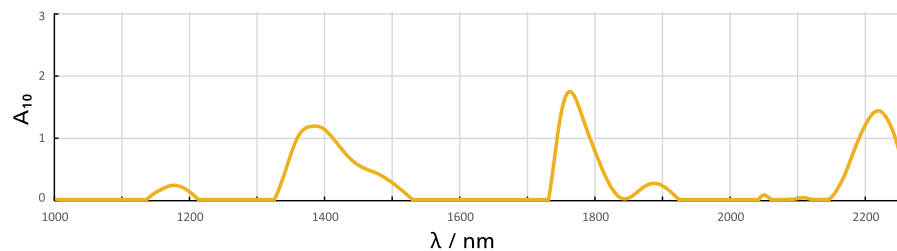


Abbildung 27 Beispiel mit gesättigten Wellenlängenbereichen.

Weitere Gründe

Es gibt noch andere Gründe für die Einbeziehung oder den Ausschluss von Wellenlängenbereichen. Die Auswahl kann auf der Kenntnis des Analyseparameters und seiner jeweiligen Absorptionsbanden beruhen (*siehe "Licht und seine Wechselwirkung mit Materie", Kapitel 2.1, Seite 3*). Dabei ist jedoch zu berücksichtigen, dass die relevanten Informationen je nach Art der Vorbehandlung ggf. in andere Wellenlängenbereiche verschoben sein können.

Wurden bei der Datenvorbehandlung Anomalien am Anfang und am Ende der Spektren eingeführt, können die entsprechenden Wellenlängen ausgeschlossen werden.

Variationen bei chemischen Komponenten oder Schwankungen der Umgebungsbedingungen können bestimmte Wellenlängenbereiche beeinflussen. Der Ausschluss dieser Wellenlängenbereiche kann die Robustheit des Modells verbessern.

i Beim Ausschluss wohlgeformter Wellenlängenbereiche ist Vorsicht geboten. Bereiche, die scheinbar keine Informationen enthalten, können in Wirklichkeit versteckte und wichtige Informationen liefern. Sie können hilfreich bei der Erkennung von Ausreißern oder bei der Verarbeitung von interferierenden Absorptionsbanden sein. Tatsächlich sind interferierende Absorptionsbanden der Hauptgrund für die Durchführung multivariater Messungen (*siehe "Beispiel für eine lineare Regression", Kapitel 6.1, Seite 91*).

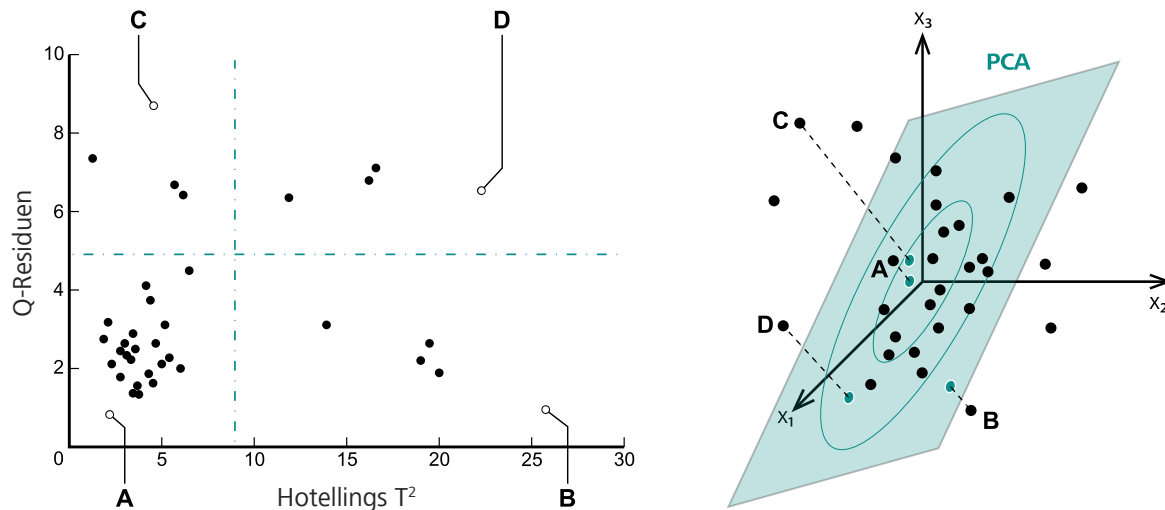


Abbildung 28 Influence-Plot (links), ursprünglicher Raum und Raum latenter Variablen (rechts). Jedes Spektrum wird durch einen Punkt in der linken Abbildung und einen Punkt in der rechten Abbildung dargestellt.

In beiden Ansichten sind 4 Punkte mit unterschiedlichen Merkmalen hervorgehoben:

- Spektrum A hat niedrige Scores und niedrige Residuen. Es liegt nahe am Modellmittelpunkt und wird vom Modell gut erklärt.
- Spektrum B ist ein Hotellings T^2 -Ausreisser. Es liegt abseits des Mittelpunkts, wird aber vom Modell gut erklärt.
- Spektrum C ist ein Q-Residuen-Ausreisser. Es liegt abseits des Mittelpunkts und wird vom Modell schlecht erklärt.
- Das Spektrum D ist sowohl ein Hotellings T^2 -Ausreisser als auch ein Q-Residuen-Ausreisser. Es liegt abseits des Mittelpunkts und wird vom Modell nur teilweise erklärt.

Der Influence-Plot zeigt, wie verschiedene Spektren das Modell beeinflussen. Da alle latenten Variablen durch den Mittelpunkt gehen, haben Spektren in der Nähe des Mittelpunkts (z. B. Spektrum A) kaum eine Chance, die Richtung der latenten Variablen zu ändern. Sie haben keinen Hebelwert. Je weiter die Distanz zum Mittelpunkt, umso grösser ist der Hebelwert und das Potenzial zur Beeinflussung des Modells. Einigen Spektren gelingt es tatsächlich, das Modell in ihre Richtung zu ziehen (Spektrum B), während das bei anderen nur in gewissem Masse (Spektrum D) oder gar nicht der Fall ist (Spektrum C).

Verglichen mit einem Modell auf Basis aller Spektren wird die Berechnung eines Modells ohne Spektrum B das Modell wahrscheinlich stärker verändern als eines ohne Spektrum D und noch stärker als ohne Spektrum C. Spektrum B beeinflusst das Quantifizierungsmodell wahrscheinlich stark – zum Guten oder zum Schlechten. Die Entscheidung, ob ein potenzieller

4.3.4 Referenzwert-Ausreisser (Quantifizierung)

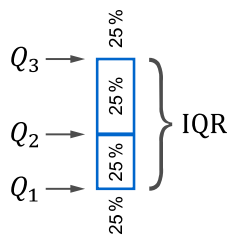
Bei Quantifizierungsmodellen werden neben den spektralen Ausreissern auch Referenzwert-Ausreisser ermittelt. Referenzwert-Ausreisser zeigen Anomalien im Referenzwert an.

Typischerweise handelt es sich bei Referenzwert-Ausreissern um falsch übermittelte Zahlen, z. B. 143 statt 14.3 oder 15.9 statt 51.9. Die Ausreissererkennung identifiziert solche Übertragungs- oder Transkriptionsfehler auf der Grundlage eines empirischen Ansatzes. Nur offensichtliche Fehler werden zur weiteren Untersuchung gekennzeichnet.

Box-Plots

Referenzwert-Ausreisser werden mithilfe einer Methode auf der Basis von Box-Plots ermittelt. Ein **Box-Plot** sortiert die Referenzwerte in aufsteigender Reihenfolge. Quartile unterteilen den Datensatz in 4 Teile. Jeder Teil enthält 25 % der Referenzwerte.

Das erste Quartil Q_1 trennt die 25 % niedrigsten Werte vom Rest. Q_2 ist der Median und trennt die 50 % niedrigsten Werte vom Rest. Das dritte Quartil Q_3 trennt die 75 % niedrigsten Werte vom Rest. Ein vertikales Rechteck stellt die mittleren 50 % der Daten dar, den Interquartilsabstand (IQR, englisch interquartile range).



Daten, die um einen bestimmten Betrag ausserhalb der IQR-Box liegen, werden als potenzielle Ausreisser betrachtet und können als kleine Kreise dargestellt werden. Der untere und obere Grenzwert für Ausreisser wird häufig als das 1.5-fache des IQR definiert:

$$[Q_1 - 1.5 \text{ IQR} ; Q_3 + 1.5 \text{ IQR}]$$

Hierbei entspricht Q_1 dem ersten Quartil, Q_3 dem dritten Quartil und IQR dem Interquartilsabstand ($Q_3 - Q_1$).

Zur Vervollständigung des Box-Plots reichen die Whisker oberhalb und unterhalb der Box bis zu den am weitesten entfernten Punkten, die nicht als potenzielle Ausreisser gekennzeichnet sind.

Anpassung für schiefe Verteilungen

Der übliche Box-Plot geht von einer annähernd symmetrischen Verteilung der Daten aus. Bei schiefen Verteilungen werden in der Regel viele reguläre Referenzwerte als potenzielle Ausreisser markiert. Aus diesem Grund nutzt die OMNIS Software eine angepasste Version des Box-Plots, welche die Schiefe der Verteilung berücksichtigt.

Im folgenden Beispiel ist die Verteilung zu den höheren Referenzwerten hin verschoben. Daraus ergeben sich 8 Ausreisser mit hohen Werten. Nach Anpassung der Ausreissergrenzwerte bleiben nur noch 2 von ihnen übrig. Auf der anderen Seite erscheint ein neuer Ausreisser mit niedrigem Wert.

4.4.2 Validierung von Quantifizierungsmodellen

Bei der Validierung wird geprüft, ob das Modell die Anforderungen an seine Leistungsfähigkeit und Robustheit erfüllt. Dazu muss der zu erwartende Fehler der Prädiktion so realistisch wie möglich abgeschätzt werden.

Ein Quantifizierungsmodell wird in der Regel wie folgt validiert:

1. Ausgehend von einer begrenzten Probenanzahl und ohne einen Validierdatensatz wird ein Modell entwickelt und mittels Kreuzvalidierung getestet (siehe unten).
2. Mit einer grösseren Probenanzahl wird der Datensatz in einen Kalibrierdatensatz zur Entwicklung eines Modells und einen Validierdatensatz zur Validierung des Modells aufgeteilt.
3. Schliesslich werden Proben für den Validierdatensatz an einem anderen Tag gesammelt und gemessen, ggf. von einer anderen Person und mit einem anderen Gerät.

Kreuzvalidierung

Die Kreuzvalidierung stützt sich ausschliesslich auf den Kalibrierdatensatz. Sie ergibt einen Schätzwert für jede Kalibrierprobe unter Verwendung eines temporären Modells, das ohne die jeweilige Kalibrierprobe erstellt wurde.

Bei der Kreuzvalidierung wird ein Mehrundenverfahren auf eine der folgenden Arten verwendet:

▪ Leave-One-Out

In jeder Runde wird bei der Leave-One-Out-Kreuzvalidierung (LOO-CV) 1 Probe zurückgestellt, während aus den übrigen Proben ein Modell erstellt wird. Dieses Modell sagt dann den Analyseparameter für die zurückgestellte Probe vorher. Diese Prädiktion dient als Schätzwert für die Probe.

Der Zyklus wird fortgesetzt, bis jede Probe einmal zurückgestellt wurde.

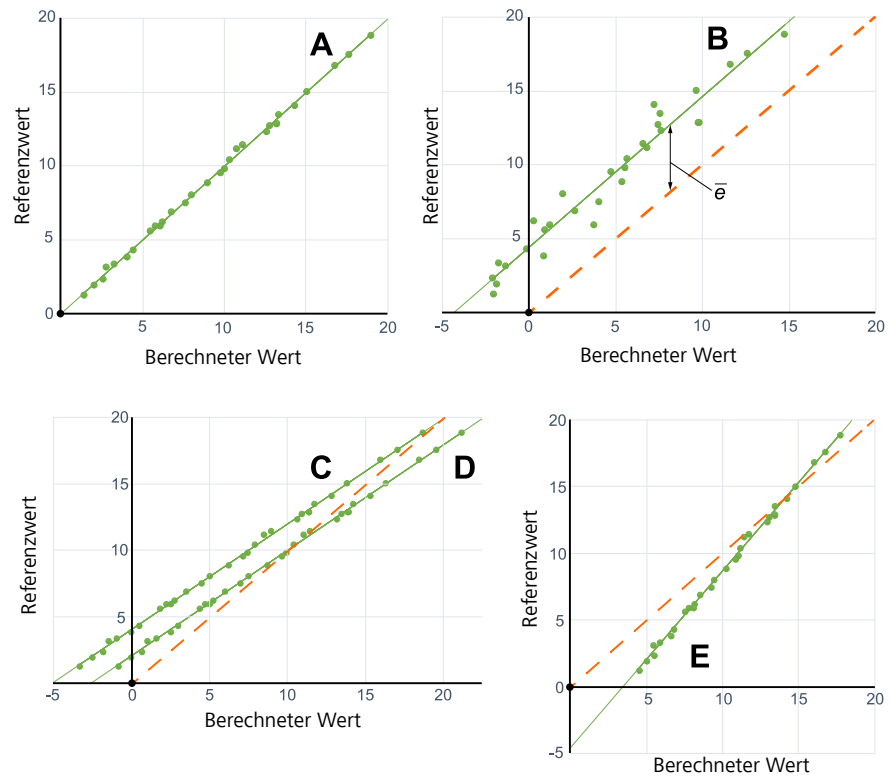


Abbildung 32 Korrelationsdiagramme. Jeder Punkt steht für eine Probe und zeigt ihren Referenzwert sowie ihren berechneten Wert. Die gestrichelte Linie zeigt die ideale 45°-Gerade.

Systematische Fehler

Systematische Fehler sind Fehler, die immer auftreten und für eine bestimmte Anwendung wiederholbar sind. Systematische Fehler können korrigiert werden. Sie werden durch den Bias $\bar{e}$ und die Steigung b der Regressionsgeraden quantifiziert:

$$y = b\hat{y} + \bar{e}$$

Falls die Steigung gleich 1 und der Bias gleich 0 ist, dann liegen keine systematischen Fehler vor.

Der **Bias** ist der durchschnittliche Fehler zwischen den Referenzwerten und den berechneten Werten:

$$\bar{e} = \frac{1}{n} \sum_{i=1}^n e_i = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i) = \bar{y} - \bar{\hat{y}}$$

Hierbei entspricht n der Probenanzahl, e_i dem Fehler der i -ten Probe, y_i dem Referenzwert der i -ten Probe, $\hat{y}_i$ dem berechneten Wert der i -ten Probe, $\bar{y}$ dem Mittelwert der Referenzwerte und $\bar{\hat{y}}$ dem Mittelwert der berechneten Werte.

Systematische Fehler				Zufällige Fehler
Gerade	Bias	Steigung	y-Achsenabschnitt	
D	~ 0	< 1	> 0	klein
E	< 0	> 1	< 0	klein

4.4.2.2 Statistische Kennzahlen

Statistische Kennzahlen drücken die Übereinstimmung zwischen Referenzwerten und vorhergesagten Werten in Zahlen aus. Die vorhergesagten Werte werden durch das Quantifizierungsmodell ermittelt.

R² – Bestimmtheitsmass

Das **Bestimmtheitsmass R²** (englisch coefficient of determination) misst die Anpassungsgüte des Quantifizierungsmodells. Für einen bestimmten Datensatz ist dies der durch das Quantifizierungsmodell erklärte Anteil der Referenzwertvariation:

$$R^2 = \frac{SS_{\text{reg}}}{SS_{\text{tot}}} = 1 - \frac{SS_{\text{res}}}{SS_{\text{tot}}} = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2}$$

Hierbei entspricht SS_{reg} der Regressionsquadratsumme (Varianz der vorhergesagten Werte, d. h. erklärte Varianz), SS_{tot} der totalen Quadratsumme (Referenzwertvarianz), SS_{res} der Residuenquadratsumme (Residualvarianz, d. h. unerklärte Varianz), y_i dem Referenzwert der i -ten Probe, $\hat{y}_i$ dem vorhergesagten Wert der i -ten Probe, und $\bar{y}$ dem Mittelwert der Referenzwerte.

Der R²-Wert ist ein Bruchteil von 1. Ein R² von 1 bedeutet, dass die vorhergesagten Werte perfekt mit den Referenzwerten übereinstimmen. Ein R² von 0.9 bedeutet, dass 90 % der Referenzwertvarianz durch die vorhergesagten Werte erklärt werden und 10 % nicht erklärt werden.

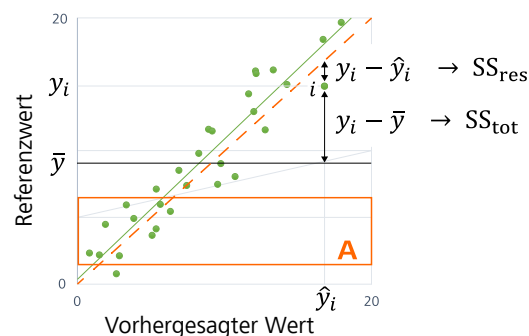


Abbildung 33 Die Komponenten zur Berechnung von R². Die gestrichelte Linie ist die 45°-Gerade.

$$R^2P = r_{\bar{v}, \hat{v}}^2 = \frac{(\sum_{i=1}^v (v_i - \bar{v}) (\hat{v}_i - \bar{\hat{v}}))^2}{\sum_{i=1}^v (v_i - \bar{v})^2 \cdot \sum_{i=1}^v (\hat{v}_i - \bar{\hat{v}})^2}$$

Hierbei entspricht v_i dem Referenzwert der i -ten Validierprobe, $\bar{v}$ dem Mittelwert der Referenzwerte, $\hat{v}_i$ dem vorhergesagten Wert der i -ten Validierprobe, $\bar{\hat{v}}$ dem Mittelwert der vorhergesagten Werte, und v der Anzahl der Validierproben.

SEC – Standardfehler der Kalibrierung

Der **Standardfehler der Kalibrierung (SEC)** basiert auf dem Kalibrierdatensatz. Der SEC kann als Schätzwert für die theoretisch beste Prädiktionsgenauigkeit angesehen werden. SEC ist die Standardabweichung der Residuen der Partial-Least-Squares-Regression (PLS):

$$SEC = \sqrt{\frac{\mathbf{e}^t \mathbf{e}}{n - k - 1}}$$

Hierbei entspricht $\mathbf{e}$ dem Residuenvektor, der alle Referenzwertvariationen im Kalibrierdatensatz enthält, die nicht durch das Modell beschrieben werden, n der Anzahl Kalibrierproben, und k der Anzahl latenter Variablen. Der Nenner $n-k-1$ ist die Anzahl der Freiheitsgrade des Residuenvektors $\mathbf{e}$.

Mit anderen Worten: Der SEC ist die Standardabweichung der Differenzen zwischen den Referenzwerten und den vorhergesagten Werten für die Proben im Kalibrierdatensatz:

$$SEC = \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n - k - 1}}$$

Hierbei entspricht y_i dem Referenzwert der i -ten Kalibrierprobe, $\hat{y}_i$ dem vorhergesagten Wert der i -ten Kalibrierprobe, n der Anzahl Kalibrierproben, und k der Anzahl latenter Variablen.

Der SEC wird manchmal auch als RMSEC bezeichnet. Der SEC enthält die zufälligen Fehler und die systematischen Fehler (Steigung und Bias).

SECV – Standardfehler der Kreuzvalidierung

Der **Standardfehler der Kreuzvalidierung (SECV)** basiert auf dem Kalibrierdatensatz. Der SECV schätzt die Prädiktionsgenauigkeit auf der Grundlage des Kalibrierdatensatzes und eines Kreuzvalidierungsverfahrens (siehe "Kreuzvalidierung", Seite 65). Der SECV kann für eine erste Modellbewertung oder zur Ermittlung der optimalen Anzahl latenter Variablen verwendet werden.

Der SECV ist die Standardabweichung der Differenzen zwischen den Referenzwerten und den Schätzwerten der Kreuzvalidierung für die Proben im Kalibrierdatensatz.

- Änderungen im Referenzmessverfahren, z. B. neues Labor, neue Ausrüstung.
- Änderungen der Proben, z. B. bei der Handhabung, Lagerung oder dem Transport.

Die Biaskorrektur und noch mehr die Steigungs-/y-Achsenabschnittskorrektur sollten mit Bedacht angewendet werden.

Falls die systematischen Fehler nicht signifikant sind, sollte keine Korrektur angewendet werden. Falls die Fehler signifikant sind, sollten sie gründlich untersucht werden. Die Ursache der Fehler sollte nach Möglichkeit behoben werden. Können die Fehler aus einem berechtigten Grund nicht behoben werden, kann eine Biaskorrektur oder eine Steigungs-/y-Achsenabschnittskorrektur angewendet werden.

Für einen verlässlichen Schätzwert des Bias werden mindestens 20 Proben benötigt. Für einen verlässlichen Schätzwert der Steigung werden mindestens 30 Proben benötigt.

Biaskorrektur

Das folgende Korrelationsdiagramm zeigt eine Biaskorrektur. Die Steigung der ursprünglichen Regressionsgeraden **F** bleibt unverändert. Nach der Korrektur (Regressionsgerade **G**) heben sich die positiven Fehler und die negativen Fehler gegenseitig auf.

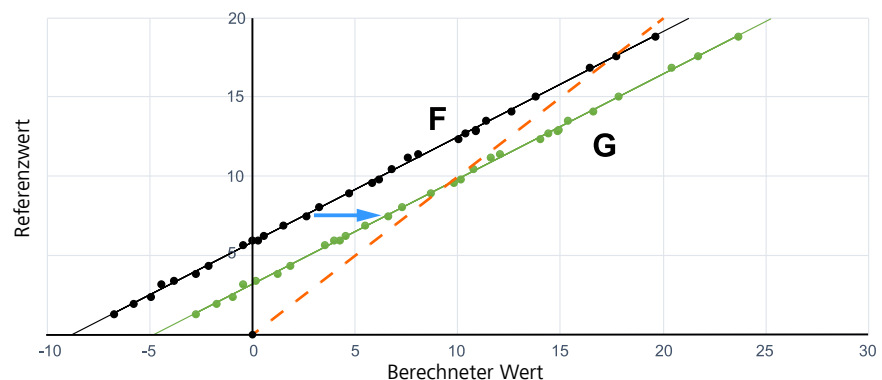


Abbildung 34 Biaskorrektur

Steigungs-/y-Achsenabschnittskorrektur

Das folgende Korrelationsdiagramm zeigt eine Steigungs-/y-Achsenabschnittskorrektur. Die ursprüngliche Regressionsgerade **H** wird um Steigung und y-Achsenabschnitt korrigiert. Infolgedessen werden sowohl die Steigung als auch der Bias korrigiert (Regressionsgerade **K**).

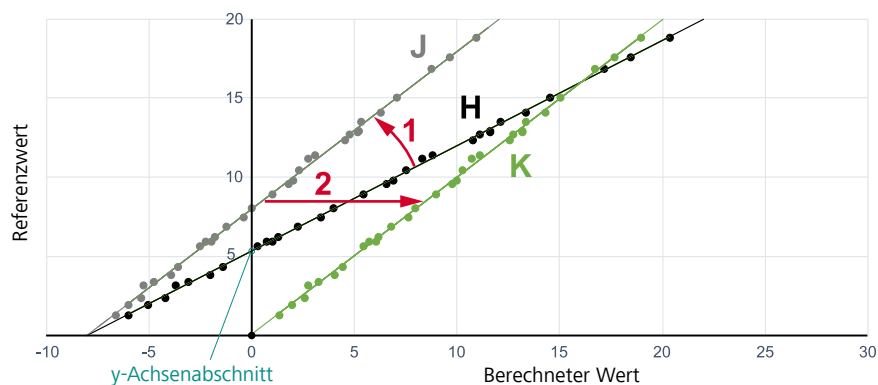


Abbildung 35 Steigungs-/y-Achsenabschnittskorrektur

SEP

Die Proben im Korrekturdatensatz bilden die Grundlage für die Steigungs-/y-Achsenabschnittskorrektur. Auf Basis dieser Proben berechnet die OMNIS Software die folgenden **Standardfehler der Prädiktion (SEP)**. Die Nenner berücksichtigen jeweils die entsprechenden Freiheitsgrade:

Typ der Korrektur	SEP
Unkorrigiert	$\text{SEP} = \sqrt{\frac{\sum_{i=1}^n (v_i - \hat{v}_i)^2}{n}}$
Biaskorrektur	$\text{SEP} = \sqrt{\frac{\sum_{i=1}^n (v_i - \hat{v}_i)^2}{n - 1}}$
Steigungs-/y-Achsenabschnittskorrektur	$\text{SEP} = \sqrt{\frac{\sum_{i=1}^n (v_i - \hat{v}_i)^2}{n - 2}}$

Hierbei entspricht v_i dem Referenzwert der i -ten Probe im Korrekturdatensatz, $\hat{v}_i$ dem vorhergesagten Wert der i -ten Probe im Korrekturdatensatz, und n der Anzahl der Proben im Korrekturdatensatz.

i Der SEP enthält alle Fehler: die zufälligen Fehler und die systematischen Fehler (Steigung und Bias).

4.5 Identifizierung und Verifizierung

4.5.1 Support Vector Machine (SVM)

Identifizierungsmodelle (ab OMNIS Software-Version 4.0) verwenden Support Vector Machines (SVM) für die Klassifizierung zwischen verschiedenen Produkten. Eine Support Vector Machine ist ein überwachter maschineller Lernalgorithmus. Anhand der Kalibrierproben lernt sie, neue Proben einem Produkt zuzuordnen.

i Der Einfachheit halber wird im Folgenden die Klassifizierung zwischen 2 Produkten beschrieben. Das Konzept ist erweiterbar, das endgültige Modell klassifiziert zwischen einer beliebigen Anzahl von Produkten.

Lineare Klassifizierung

Abbildung 36 (links) zeigt Eingangsdaten mit 2 Variablen.

i Der Einfachheit halber werden die parametrisierten Spektren in einem 2-dimensionalen Variablenraum dargestellt. Jeder Punkt stellt ein Spektrum dar, die Farbe gibt die Produktzugehörigkeit an.

Die Produkte sind linear trennbar. Der SVM-Algorithmus erstellt eine Hyperebene zwischen den Produkten (Abbildung rechts).

i Hyperebenen sind eine Verallgemeinerung von Ebenen aus dem 3-dimensionalen Raum auf Räume beliebiger Dimension. Die Dimension einer Hyperebene ist um eins kleiner als die Dimension des sie umgebenden Raums. Eine Hyperebene in einem 2-dimensionalen Raum ist eine Linie.

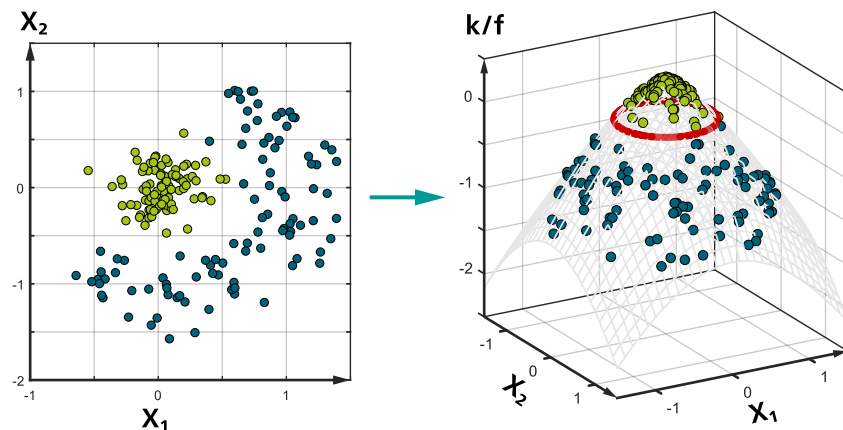


Abbildung 37 Nicht linear trennbare Produkte (links). Eine zusätzliche Dimension k/f (Kernel-Merkmal) erleichtert die Trennung (rechts). Die Werte sind in willkürlichen Einheiten angegeben.

In der Abbildung werden die Daten vom 2-dimensionalen Raum in den 3-dimensionalen Raum umgewandelt. Die Punkte des einen Produkts werden über die ursprüngliche Ebene angehoben, während die Punkte des anderen Produkts nach unten verschoben werden. Die lineare Hyperebene zwischen den Produkten ist eine 2-dimensionale Ebene, die der ursprünglichen Ebene sehr ähnlich ist. In 2 Dimensionen betrachtet, ist die Entscheidungsgrenze eine nichtlineare Linie, in diesem Fall die rote, kreisförmige Linie.

Auch hier dient die Hyperebene als Klassifikator. Eine neue Probe kann einem Produkt zugeordnet werden, je nachdem, auf welche Seite der Hyperebene ihr Spektrum fällt.

Die OMNIS Software nutzt einen Kernel mit radialer Basisfunktion, um die Daten in den Merkmalsraum zu transformieren. Der Kernel verwendet einen Skalierungsparameter, der den Grad der Nichtlinearität steuert.

Parameterauswahl

Für den Regularisierungsparameter, der die Position der Hyperebene steuert, und für den Skalierungsparameter, der den Grad der Nichtlinearität steuert, müssen geeignete Werte gewählt werden.

Beide Parameter wirken sich auf die Generalisierungsfähigkeit der Support Vector Machine aus, d. h. darauf, wie gut sie von den Kalibrierspektren auf neue, unbekannte Spektren verallgemeinern kann. Ermöglicht zum Beispiel der Skalierungsparameter einen hohen Grad an Nichtlinearität, ist die Hyperebene möglicherweise zu stark an die Kalibrierspektren angepasst (Überanpassung). Ermöglicht der Skalierungsparameter nur einen geringen Grad an Nichtlinearität, ist die Hyperebene möglicherweise nicht ausreichend an die Kalibrierspektren angepasst (Unteranpassung).

1. Positiver Validierdatensatz:
 - a. Bei einer begrenzten Anzahl positiver Proben wird zunächst kein positiver Validierdatensatz erstellt.
 - b. Sobald eine ausreichende Anzahl positiver Proben zur Verfügung steht, wird mit Hilfe der automatischen Datensatzaufteilung ein positiver Validierdatensatz erstellt.

Negativer Validierdatensatz:

- a. Die gesammelten negativen Proben werden dem negativen Validierdatensatz zugeordnet.
 - b. Die automatische Ermittlung von negativen Spektren kann eingesetzt werden, um spektrale Ausreisser zu erkennen und dem negativen Validierdatensatz zuzuordnen (*siehe "Spektrale Ausreisser – Algorithmus", Kapitel 6.5, Seite 99*).
2. Schliesslich werden Proben für den positiven und den negativen Validierdatensatz an einem anderen Tag gesammelt und gemessen, ggf. von einer anderen Person und mit einem anderen Gerät.

i Die Proben im positiven und im negativen Validierdatensatz werden für die Berechnung des Modells nicht verwendet.

Validierung

Für jede Probe ermittelt das Qualifizierungsmodell ein Resultat (positiv oder negativ). Für die Proben im Kalibrierdatensatz und im positiven Validierdatensatz wird ein positives Resultat erwartet. Für die Proben im negativen Validierdatensatz wird ein negatives Resultat erwartet. Falls das Resultat der Erwartung entspricht, ist die Prädiktion richtig (= erfolgreich), andernfalls falsch (= fehlgeschlagen).

Die OMNIS Software zeigt für Qualifizierungsmodelle die folgenden Werte an:

Erfolgreich % (Gesamt) misst die Gesamtkorrektheit eines Modells.

$$\text{Erfolgreich \% (Gesamt)} = \frac{\text{korrekte Prädiktionen}}{\text{alle Prädiktionen}}$$

Ähnliche Zahlen für *Erfolgreich %* sind für jeden Datensatz verfügbar. Im Idealfall liegen alle Kennzahlen bei 100 %.

Qualifizierung verbessern

Eine geeignetere Parametrierung (Datenvorbehandlung und Wellenlängenauswahl) kann das Qualifizierungsmodell verbessern.

Nicht qualifizierte Proben

Falls Proben fälschlicherweise nicht qualifiziert werden:

- Proben und Probenhandhabung auf Anomalien prüfen.
- Datenvorbehandlung und Wellenlängenauswahl prüfen.

- Die Spektren der nicht qualifizierten Proben im Score Plot prüfen:
 - Falls die Spektren nicht bereits im Modell enthalten sind: Die Spektren zum positiven Validierdatensatz hinzufügen.
 - Im Score Plot die Scores der zu prüfenden Spektren mit den Scores der Spektren im Kalibrierdatensatz vergleichen.

Falls die Proben keine Outlier sind, aber ihre Variationen im Kalibrierdatensatz unterrepräsentiert sind, sollte der Kalibrierdatensatz entsprechend erweitert werden.

5 Prädiktion

5.1 Quantifizierung

Bei der Prädiktion einer Probe mit unbekanntem Analyseparameter wird wie folgt vorgegangen:

1. Das Spektrum der Probe wird aufgenommen.
2. Das Quantifizierungsmodell wendet dieselbe Datenvorbehandlung und dieselben Wellenlängenbereiche wie für die Spektren im Kalibrierdatensatz an.
3. Anhand des resultierenden Spektrums sagt das Modell den Analyseparameter vorher.

i Bei der Berechnung des Quantifizierungsmodells wurden die Kalibrierspektren und die Referenzwerte mittelwertzentriert. Das wird natürlich beim oben beschriebenen Vorgehen berücksichtigt.

5.1.1 Ausreisser und Resultatüberwachung

Bei der Prädiktion eines quantitativen Analyseparameters können unterschiedliche Arten von Ausreissern erkannt und die Resultate überwacht werden:

Ausreisser/Überwachung		Ursache (Beispiel)
Spektrale Ausreisser	Hotellings T^2	Die Konzentration einer chemischen Komponente liegt ausserhalb des Konzentrationsbereichs der Kalibrierproben.
	Q-Residuen	Die Probe enthält Bestandteile, die in den Kalibrierproben nicht vorhanden sind.
Nearest Neighbor-Ausreisser (optional, ab OMNIS Software-Version 4.2)		Die Probe weist Probenvariationen auf, die in dieser Kombination in den Kalibrierproben nicht vertreten sind.
Resultatüberwachung (optional)		Der Resultatwert des Analyseparameters liegt ausserhalb der vom Kunden gewählten Akzeptanzkriterien.

Spektrale Ausreisser

Spektrale Ausreisser werden wie folgt ermittelt:

1. Die Software berechnet die Werte für Hotellings T^2 und Q-Residuen für das Spektrum auf Basis des Quantifizierungsmodells (PLS-Modell).
2. Falls T^2 oder der Q-Residuenwert des Spektrums grösser als der entsprechende vom Modell berechnete kritische Wert ist, wird die Probe in Bezug auf das angewendete Modell als Ausreisser gekennzeichnet (*siehe "Ausreisserbewertung bei der Prädiktion (Quantifizierung)", Seite 101*).

 Die T^2 und Q-Residuenwerte sind in der OMNIS Software als Variable verfügbar. Sie können mit den Werten im PLS-Influence-Plot des Quantifizierungsmodells verglichen werden. Die gestrichelten Linien geben die kritischen Werte an.

Nearest Neighbor-Ausreisser

(ab OMNIS Software-Version 4.2)

Im Idealfall decken die Kalibrierproben alle möglichen Kombinationen der Probenvariationen ab. In der Realität treten einige Kombinationen häufiger auf, andere überhaupt nicht. Dementsprechend sind die Kalibrierproben im Raum der latenten Variablen ungleich verteilt. In manchen Bereichen finden sich viele Kalibrierproben, dazwischen gibt es Lücken.

Falls das Spektrum einer unbekannten Probe in eine Lücke zwischen den Kalibrierproben fällt, kann das Prädiktionsresultat ungültig oder ungenau sein. Um solche Fälle zu erkennen, wird die Distanz D von der unbekannten Probe i zu jeder Kalibrierprobe u berechnet:

$$D = \sqrt{(\mathbf{s}_i - \mathbf{s}_u)^t (\mathbf{s}_i - \mathbf{s}_u)}$$

Hierbei entspricht $\mathbf{s}_i$ den Scores der unbekannten Probe i und $\mathbf{s}_u$ den Scores der Kalibrierprobe u . Die Scores sind normalisiert und orthogonal.

Die kleinste Distanz ist die Distanz zur nächstgelegenen Kalibrierprobe und wird als **Nearest Neighbor Distance** (NND) bezeichnet:

Falls der NND-Wert einen bestimmten NND-Grenzwert übersteigt, wird die unbekannte Probe als Nearest Neighbor-Ausreisser bezeichnet.

Der NND-Grenzwert wird wie folgt bestimmt:

1. Für jede Kalibrierprobe wird ein NND-Wert bestimmt. Dieser Wert entspricht der Distanz zur nächstgelegenen der übrigen Kalibrierproben.
2. Der maximale NND-Wert aller Kalibrierproben ist der NND-Grenzwert.

Der NND-Wert der unbekannten Probe und der NND-Grenzwert sind in der OMNIS Software als Variable verfügbar.

Resultatüberwachung

In der OMNIS Software kann die Resultatüberwachung verwendet werden, um Warngrenzen und Eingreifgrenzen für den Bereich der Prädikti-

onsresultate zu definieren. Optional können Aktionen definiert werden, die bei einer Verletzung der definierten Grenzen ausgelöst werden.

5.2 Identifizierung und Verifizierung

Bei der Identifizierung oder Verifizierung einer Probe wird wie folgt vorgegangen:

1. Das Probenspektrum wird aufgenommen.
2. Das Identifizierungsmodell wendet dieselbe Datenvorbehandlung und dieselbe Wellenlängenauswahl wie für die Spektren im Kalibrierdatensatz an.
3. Das Identifizierungsmodell wertet das Spektrum aus (*siehe "Prädiktion der Produktzugehörigkeit einer Probe", Kapitel 4.5.2, Seite 80*).
4. Falls das Identifizierungsmodell Teil einer Modellhierarchie ist und ein weiteres Identifizierungsmodell mit dem ermittelten Produkt verknüpft ist, wird dieses Modell ausgeführt. Falls ein oder mehrere Quantifizierungsmodelle mit dem ermittelten Produkt verknüpft sind, werden diese Modelle ausgeführt.
5. Das Identifizierungsergebnis oder das Verifizierungsergebnis wird angezeigt, bei einer Modellhierarchie ggf. auch Quantifizierungsergebnisse.

Identifizierungsstatus

- Identifiziert
Die Identifizierung ist erfolgreich.
- Mehrdeutig
Mehrere Produkte übertreffen die Wahrscheinlichkeitsschwelle. Die Identifizierung ist fehlgeschlagen.
- Nicht identifiziert
Kein Produkt übertrifft die Wahrscheinlichkeitsschwelle. Die Identifizierung ist fehlgeschlagen.

Verifizierungsstatus

- Erfolgreich
Die Probe wurde erfolgreich identifiziert und das Ergebnis entspricht dem erwarteten Produkt.
- Fehlgeschlagen
Die Verifizierung ist fehlgeschlagen.

5.3 Qualifizierung

Bei der Qualifizierung einer Probe wird wie folgt vorgegangen:

1. Das Probenspektrum wird aufgenommen.
2. Das Qualifizierungsmodell wendet dieselbe Datenvorbehandlung und dieselbe Wellenlängenauswahl wie für die Spektren im Kalibrierdatensatz an.
3. Anhand des resultierenden Spektrums qualifiziert das Modell die Probe.
4. Das Qualifizierungsergebnis wird angezeigt.

Qualifizierungsstatus

- Erfolgreich
- Fehlgeschlagen

6 Anhang

6.1 Beispiel für eine lineare Regression

Univariate lineare Regression

Im einfachsten Fall hat ein Gemisch nur einen Absorber und das Spektrum nur einen Peak. Proben mit unterschiedlichen Konzentrationen des Absorbers weisen Peaks mit unterschiedlichen Absorbanzwerten auf (*siehe "Beer-Lambert-Gesetz", Kapitel 2.2.1, Seite 6*).

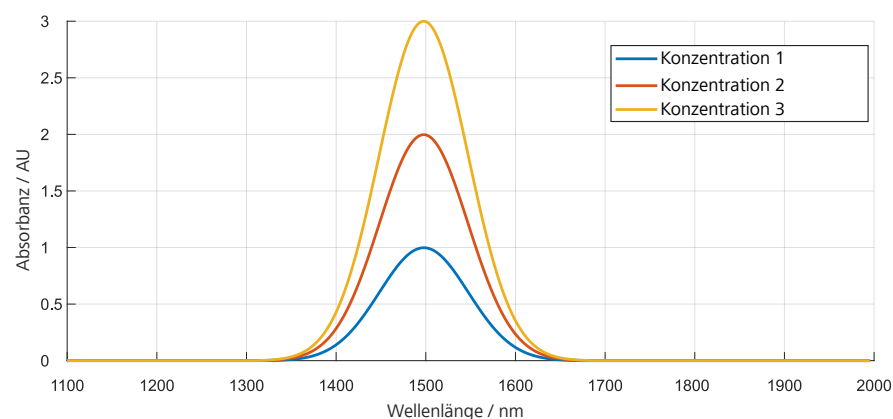


Abbildung 38 Modellierte Daten für 3 Proben mit einem Peak bei 1'500 nm.

Die 3 gemessenen Absorbanzwerte bei 1'500 nm können gegen die Konzentration des Absorbers aufgetragen werden.

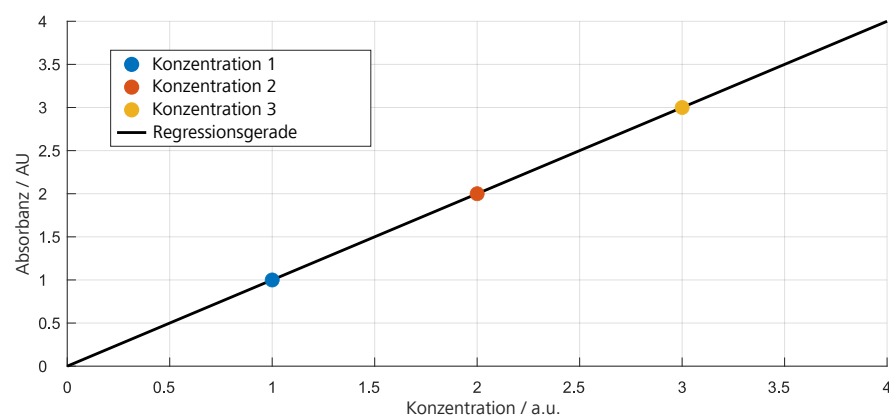


Abbildung 39 Beziehung zwischen Absorbanzwerten und Konzentrationen.

Gemäss dem Beer-Lambert-Gesetz ist die Beziehung zwischen den Absorbanzwerten und den Konzentrationen linear. Mit dem Satz von

3 Proben kann eine lineare Regression durchgeführt werden, die eine Regressionsgerade wie abgebildet ergibt.

Da es nur 1 Variable gibt (1 Wellenlänge), handelt es sich um eine univariate lineare Regression. Diese Regression kann als Quantifizierungsmodell verwendet werden. Bei einer Probe mit unbekannter Konzentration wird die Absorbanz A bei 1'500 nm gemessen. Aus der Regressionsgeraden ergibt sich dann die entsprechende Konzentration c des Absorbers:

$$c = bA$$

Der Koeffizient b ist konstant und identisch mit der Steigung der Regressionsgeraden.

Es ist zu beachten, dass alle Proben den gleichen Absorber mit dem gleichen molaren Extinktionskoeffizienten enthalten müssen. Zudem müssen alle Absorptionsmessungen mit identischen Schichtdicken durchgeführt werden.

Multivariate lineare Regression

Reale Gemische enthalten mehr als einen Absorber. Das aufgenommene Spektrum ist die Summe aller Absorberspektren (*siehe "Beer-Lambert-Gesetz", Kapitel 2.2.1, Seite 6*).

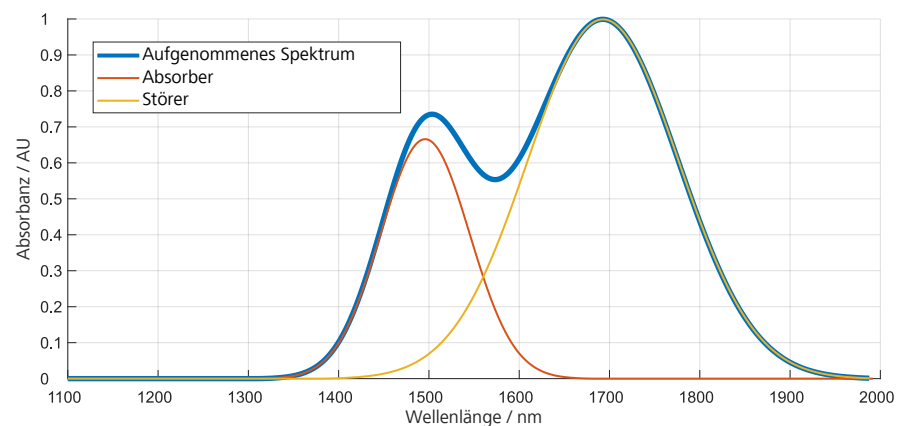


Abbildung 40 Modellierte Daten mit 2 Komponenten. Der Absorber (rote Linie) soll quantifiziert werden.

Das aufgenommene Spektrum (blaue Linie) ist die Summe aus dem reinen Absorberspektrum und einem überlappenden Störerspektrum. Bei 1'500 nm setzt sich der gemessene Absorbanzwert nicht nur aus der Absorbanz des Absorbers, sondern auch aus der Absorbanz des Störers zusammen. Der Analyseparameter kann nicht mit einer einzigen Messung bei 1'500 nm quantifiziert werden. Es ist auch unmöglich zu wissen, ob ein Störer vorhanden war und ob die Messung zuverlässig ist.

Was passiert, wenn 2 Wellenlängen z. B. bei 1'500 nm und 1'700 nm gemessen werden? Die gemessene Absorbanz bei Wellenlänge 1, A_1 , ist die Summe des reinen Absorbersignals A_1^a (Index a = Absorber) und des

reinen Störersignals A_1^f (Index f = Störer). Das Gleiche gilt für die gemessene Absorbanz bei Wellenlänge 2, A_2 :

$$\begin{aligned} A_1 &= A_1^a + A_1^f = \varepsilon_1^a c_a + \varepsilon_1^f c_f \\ A_2 &= A_2^a + A_2^f = \varepsilon_2^a c_a + \varepsilon_2^f c_f \end{aligned}$$

Hierbei entsprechen ε_1^a und ε_1^f den molaren Extinktionskoeffizienten bei Wellenlänge 1 für den Absorber bzw. den Störer, und c_a und c_f den Konzentrationen für den Absorber bzw. den Störer.

In den obigen Gleichungen ist die Schichtdicke l aus dem Beer-Lambert-Gesetz ausgeschlossen. Das macht die Algebra später ein wenig einfacher. Die Schichtdicke muss natürlich bei allen Proben gleich sein. Die Absorbanzen in den Gleichungen sind daher Absorbanzen pro cm.

Die Gleichungen können in Matrixform wie folgt geschrieben werden:

$$\begin{bmatrix} A_1 \\ A_2 \end{bmatrix} = \begin{bmatrix} \varepsilon_1^a & \varepsilon_1^f \\ \varepsilon_2^a & \varepsilon_2^f \end{bmatrix} \begin{bmatrix} c_a \\ c_f \end{bmatrix}$$

Daher:

$$\begin{bmatrix} c_a \\ c_f \end{bmatrix} = \begin{bmatrix} \varepsilon_1^a & \varepsilon_1^f \\ \varepsilon_2^a & \varepsilon_2^f \end{bmatrix}^{-1} \begin{bmatrix} A_1 \\ A_2 \end{bmatrix}$$

Die Lösung für die Konzentration des Absorbers ergibt:

$$c = c_a = \frac{\varepsilon_2^f}{\varepsilon_1^a \varepsilon_2^f - \varepsilon_1^f \varepsilon_2^a} A_1 + \frac{-\varepsilon_1^f}{\varepsilon_1^a \varepsilon_2^f - \varepsilon_1^f \varepsilon_2^a} A_2$$

Daraus ergibt sich, dass die Konzentration des Absorbers auch bei Vorhandensein eines Störers berechnet werden kann, indem die Absorbanz bei zwei Wellenlängen gemessen und jede Absorbanz mit einer Konstante multipliziert wird.

Die Konstanten beziehen sich auf die molaren Extinktionskoeffizienten und können in Tabellen nachgeschlagen werden. Das geschieht in der Realität aber nie. Stattdessen werden sie durch einen Kalibrierschritt und die Lösung des linearen Gleichungssystems durch eine multivariate lineare Regression wie die PLS-Regression ermittelt. Daher werden die Konstanten als Regressionskoeffizienten b_1 und b_2 bezeichnet:

$$c = b_1 A_1 + b_2 A_2$$

Mehr als 2 Absorber

Wie oben gezeigt genügt es bei 1 Absorber, die Absorbanz bei 1 Wellenlänge zu bestimmen. Mit 2 Absorbern genügt es, die Absorbanz bei 2 Wellenlängen zu bestimmen.

Das kann verallgemeinert werden. Mehr Absorber erfordern mehr Absorbanzwerte A_i bei verschiedenen Wellenlängen i . Es handelt sich immer noch um eine lineare Beziehung:

$$c = b_1A_1 + b_2A_2 + \cdots + b_nA_n$$

Entwicklung eines Quantifizierungsmodells

Bevor die obige Gleichung die Konzentration in unbekannten Proben vorhersagen kann, müssen die Koeffizienten b_1 , b_2 usw. bestimmt werden. Dazu ist ein Kalibrierschritt erforderlich. Es werden mehrere Proben mit unterschiedlichen Konzentrationen des Analyseparameters gemessen.

In Einklang mit der später für PCA und PLS verwendeten Terminologie kann c durch y und A durch x ersetzt werden. Dann sieht die obige Gleichung für jede Kalibrierprobe ausgeschrieben wie folgt aus:

$$\begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{bmatrix} = \begin{bmatrix} x_{1,1} & x_{1,2} & \cdots & x_{1,f} \\ x_{2,1} & x_{2,2} & \cdots & x_{2,f} \\ \vdots & \vdots & \ddots & \vdots \\ x_{n,1} & x_{n,2} & \cdots & x_{n,f} \end{bmatrix} \begin{bmatrix} b_1 \\ b_2 \\ \vdots \\ b_n \end{bmatrix} + \begin{bmatrix} e_1 \\ e_2 \\ \vdots \\ e_n \end{bmatrix}$$

Hierbei entspricht n der Probenanzahl, f der Anzahl der Wellenlängen, y_1 dem Referenzwert für die Probe 1, gemessen mit der Referenzmethode (z. B. Titration), $x_{1,1}$ der gemessene Absorbanz von Probe 1 bei Wellenlänge 1, und $\{x_{1,1} \dots x_{1,f}\}$ dem Spektrum von Probe 1, gemessen bei f Wellenlängen. $b_1 \dots b_n$ entsprechen den Regressionskoeffizienten, und $e_1 \dots e_n$ den Fehlertermen, die zeigen, wie gut die Regressionskoeffizienten die Messdaten modellieren.

In einer kompakteren Matrixform ergibt das:

$$\mathbf{y} = \mathbf{X}^t \mathbf{p} + \mathbf{e}$$

X ist definiert als Matrix $f \times n$. **X^t** ist die transponierte Matrix **X**, d. h. die Zeilen und Spalten werden vertauscht, um die obige Matrix $n \times f$ zu erhalten. Der Prädiktionsvektor **p** entspricht den obigen Regressionskoeffizienten **b**.

Der Prädiktionsvektor $\mathbf{p}$ ermöglicht die Prädikation des Analyseparameters für eine neue Probe auf Basis ihres Spektrums $\mathbf{x}$. Der vorhergesagte Wert $\hat{y}$ ist:

$$\hat{y} = \mathbf{x}^t \mathbf{p}$$

Die Aufgabe der multivariaten linearen Regression ist die Bestimmung von Regressionskoeffizienten, die minimale Fehlerterme ergeben. Eine multiple lineare Regression (MLR) würde jedoch eine grössere Anzahl an Kalibrierproben erfordern als die Anzahl der Wellenlängen. Ein weiteres Hindernis ist die hohe Korrelation zwischen den Variablen.

Für spektroskopische Prädiktionen können andere Methoden verwendet werden. Mit PCA kann die Datenmenge stark reduziert und die Korrelation

vollständig eliminiert werden. Mit einer PLS-Regression werden zusätzlich die Referenzwerte der Proben berücksichtigt.

6.2 PCA-Algorithmus

Die **Hauptkomponentenanalyse (PCA)**, englisch *principal component analysis*) wird für die automatische Datensatzaufteilung und die Erkennung spektraler Ausreißer verwendet (*siehe "Hauptkomponentenanalyse (PCA)", Kapitel 4.2, Seite 36*).

Für die Hauptkomponentenanalyse der parametrisierten und mittelwert-zentrierten Spektren führt die OMNIS Software eine **Singulärwertzerlegung (SVD)**, englisch *singular value decomposition*) durch. Sie zerlegt die ursprüngliche Datenmatrix **X** (die Spektren des Kalibrierdatensatzes) in 3 Matrizen und berechnet n Hauptkomponenten:

$$\mathbf{X} = \mathbf{L}\mathbf{\Sigma}\mathbf{S}^t$$

Hierbei entspricht **X** den ursprünglichen Spektrendaten (eine Matrix $f \times n$ mit f Wellenlängen und n Proben), **L** den Loadings (eine Matrix $f \times n$), **Σ** den Singulärwerten (eine Matrix $n \times n$), und **S** den Scores (eine Matrix $n \times n$). Die Gleichung kann in grafischer Form dargestellt werden (*siehe Abbildung 41, Seite 96*).

Die Loadings **L** projizieren die Achsen des Wellenlängenraums auf die Achsen des Hauptkomponentenraums. Die Spaltenvektoren von **L** sind die Hauptachsen oder Hauptrichtungen.

Die Scores **S** sind die Projektionen der ursprünglichen Spektren auf die Hauptachsen. Jeder Zeilenvektor von **S** enthält ein bestimmtes Spektrum.

Σ ist eine Diagonalmatrix der Singulärwerte σ_j . Die Singulärwerte beschreiben die durch jede Hauptkomponente erklärte Varianz. Durch Konvention sind sie so sortiert, dass $\sigma_1 > \sigma_2 > \dots > \sigma_n$.

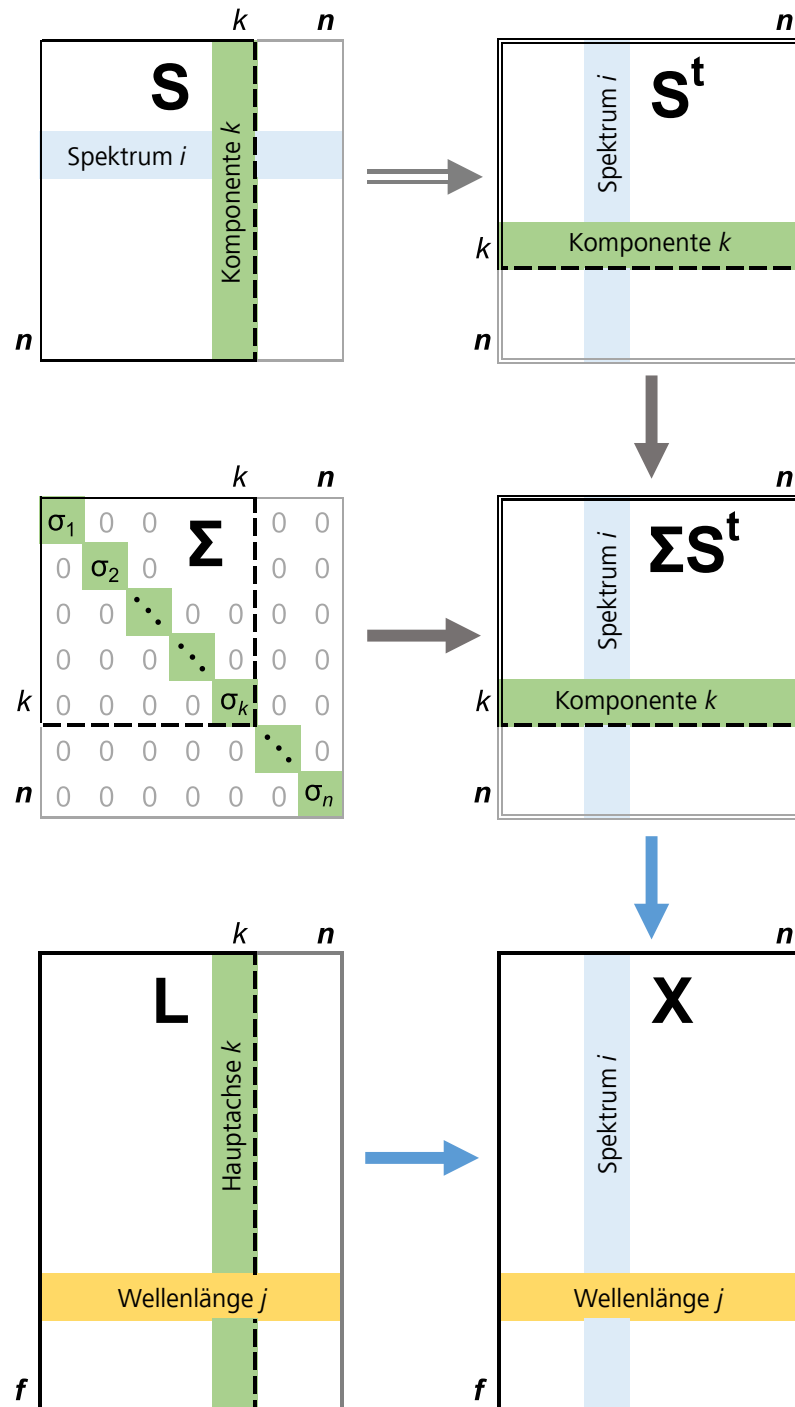


Abbildung 41 Gleichung der Singulärwertzerlegung in grafischer Form. Die gestrichelten Linien zeigen die verkürzten Matrizen für ein Modell mit k Hauptkomponenten. Die Informationen in den $n-k$ ausgelassenen Hauptkomponenten fließen in die Residualmatrix ein (eine Matrix $f \times n$, nicht abgebildet).

Residualmatrix

Ein PCA-Modell nutzt nur die ersten paar der berechneten n Hauptkomponenten. In [Abbildung 41](#) sind es die ersten k von n Hauptkomponenten. Die Originaldaten $\mathbf{X}$ können in vom Modell beschriebene Daten und nicht vom Modell beschriebene Daten unterteilt werden:

$$\mathbf{X} = \mathbf{L}_a \mathbf{\Sigma}_a \mathbf{S}_a^t + \mathbf{E}$$

Hierbei entspricht $\mathbf{X}$ den ursprünglichen Spektrendaten (eine Matrix $f \times n$), $\mathbf{L}_a$ (eine Matrix $f \times k$) den ersten k Spalten von $\mathbf{L}$, $\mathbf{\Sigma}_a$ (eine Diagonalmatrix $k \times k$) den ersten k Singulärwerten, $\mathbf{S}_a$ (eine Matrix $n \times k$ mit k Hauptkomponenten) den ersten k Spalten von $\mathbf{S}$, und $\mathbf{E}$ der Residualmatrix (eine Matrix $f \times n$), die alle spektralen Variationen in $\mathbf{X}$ enthält, die nicht vom Modell beschrieben werden können.

In der Regel ist $k \ll n \ll f$. Zum Beispiel: $k = 3$ Hauptkomponenten, $n = 100$ Proben und $f = 2'500$ Wellenlängen.

Jede Spalte $\mathbf{e}_i$ der Residualmatrix $\mathbf{E}$ zeigt die orthogonale Distanz des Spektrums i zum PCA-Raum, genannt **Residuum**. Je mehr Hauptkomponenten das Modell verwendet, umso kleiner ist das Residuum.

6.3 PLS-Algorithmus

Die **Partielle Kleinste-Quadrate-Regression (PLS-Regression)**, englisch *partial least squares regression* wird zur Berechnung des Quantifizierungsmodells verwendet ([siehe "PLS-Regression", Kapitel 4.4.1, Seite 62](#)).

PLS bezieht zwei Datenblöcke mit ein:

- Die parametrisierte und mittelwertzentrierte Matrix $\mathbf{X}$ (die Spektren).
- Den mittelwertzentrierten $\mathbf{y}$ -Vektor (die Referenzwerte).

PLS zerlegt die Matrix $\mathbf{X}$ in 2 Matrizen:

$$\mathbf{X} = \mathbf{L}\mathbf{S}^t + \mathbf{Z}$$

Hierbei entspricht $\mathbf{X}$ (eine Matrix $f \times n$ mit f Wellenlängen und n Proben) den vorbehandelten und mittelwertzentrierten Spektren, $\mathbf{L}$ den Loadings (eine Matrix $f \times k$ mit k latenten Variablen), $\mathbf{S}$ den Scores (eine Matrix $n \times k$), und $\mathbf{Z}$ der Residualmatrix (eine Matrix $f \times n$), die alle spektralen Variationen in $\mathbf{X}$ enthält, die nicht vom Modell beschrieben werden können.

Während bei PCA die Score-Matrix $\mathbf{S}$ die Varianz von $\mathbf{X}$ erklärt, erklärt bei PLS die Score-Matrix $\mathbf{S}$ die Kovarianz zwischen $\mathbf{X}$ und $\mathbf{y}$. PLS maximiert die durch die Scores erklärte Kovarianz. Das bedeutet, dass die Scores nicht nur die Variabilität von $\mathbf{X}$ am besten erklären, sondern auch die höchstmögliche Korrelation mit den Referenzwerten aufweisen.

Zur Maximierung der Kovarianz zwischen $\mathbf{X}$ und $\mathbf{y}$ tauscht der PLS-Algorithmus Daten zwischen $\mathbf{X}$ und $\mathbf{y}$ aus. Daher verschmelzen $\mathbf{X}$ und $\mathbf{y}$ zu

MD^2 kann einfach als T^2 bezeichnet werden. Ein Spektrum mit $T^2 = 0$ projiziert auf den Mittelpunkt des mittelwertzentrierten Modells, d. h. alle Scores liegen im Mittel. Je weiter sich die Projektion eines Spektrums auf die Hyperebene vom Mittelpunkt entfernt, umso höher ist der T^2 -Wert. In der Nähe des Mittelpunkts passt das Modell am besten. Weit entfernt vom Mittelpunkt passt das Modell möglicherweise schlecht.

Es kann ein **Signifikanzniveau** von z. B. 5 % für einen Hypothesentest festgelegt werden, der Ausreisser erkennt (*siehe "Spektrale Ausreisser – Algorithmus", Kapitel 6.5, Seite 99*).

Q-Residuum

Das Q-Residuum für das Spektrum i ist die quadrierte Distanz vom PCA- oder PLS-Modell zum Spektrum:

$$Q_i = \mathbf{e}_i^T \mathbf{e}_i = \sum_{j=1}^f e_{i,j}^2$$

Hierbei entspricht $\mathbf{e}_i$ der i -ten Spalte der Residualmatrix $\mathbf{E}$, $e_{i,j}$ dem Residuum für das Spektrum i , und f der Anzahl der Wellenlängen.

Die Q-Residuen zeigen die Variationen, die nicht durch das Modell erklärt werden. Ein hohes Q-Residuum zeigt an, dass das Spektrum möglicherweise nicht mit dem Modell übereinstimmt, z. B. wenn die gemessene Probe eine andere Substanz enthält.

Es kann ein **Signifikanzniveau** von z. B. 5 % für einen Hypothesentest festgelegt werden, der Ausreisser erkennt (*siehe "Spektrale Ausreisser – Algorithmus", Kapitel 6.5, Seite 99*).

6.5 Spektrale Ausreisser – Algorithmus

Die Erkennung spektraler Ausreisser identifiziert Spektren, die von der Grundgesamtheit abweichen (*siehe "Erkennung spektraler Ausreisser", Seite 53*). Der Algorithmus bewertet, ob die Werte für Hotellings T^2 oder für das Q-Residuum des geprüften Spektrums das Resultat einer zufälligen oder einer systematischen Variation sind.

3. Hotellings T^2 -Ausreisser:

Auf Basis der Hotellings T^2 -Werte (siehe "*Hotellings T^2* ", Seite 98) wird ein Hypothesentest nach H. Hotelling, *The Generalization of Student's Ratio*, The Annals of Mathematical Statistics Band 2, Nr. 3 (Aug. 1931), S. 360–378 durchgeführt.

- a. Die Nullhypothese lautet, dass der T^2 -Wert des zu untersuchenden Spektrums zu den T^2 -Werten des PCA-Modells passt. Ist die Nullhypothese wahr, folgt T^2 einer Hotellings T^2 -Verteilung. Die Verteilung kann als skalierte F -Verteilung ausgedrückt werden:

$$T^2 \sim \frac{k(n-1)}{n-k} F_{k,n-k}$$

dabei ist k die Anzahl der Hauptkomponenten, n die Anzahl der Spektren und $F_{k,n-k}$ die F -Verteilung mit den Parametern k und $n-k$.

- b. Das einstellbare Signifikanzniveau steuert die Wahrscheinlichkeit, mit der die Nullhypothese zurückgewiesen wird, falls sie wahr ist (bekannt als Typ-I-Fehler). Der Standardwert ist 5 %.
- c. Auf Basis dieser Verteilung und des Signifikanzniveaus wird ein kritischer Wert für T^2 berechnet.
- d. Falls der T^2 -Wert für das Spektrum grösser als der kritische Wert ist, wird die Nullhypothese zurückgewiesen und das Spektrum als potenzieller Ausreisser gekennzeichnet.

4. Q-Residuen-Ausreisser

Auf Basis der Q-Residuen (siehe "*Q-Residuum*", Seite 99) wird ein Hypothesentest nach J. E. Jackson und G. S. Mudholkar, *Control Procedures for Residuals Associated With Principal Component Analysis*, Technometrics Band 21, Nr. 3 (Aug. 1979), S. 341–349 durchgeführt.

- a. Die Nullhypothese lautet, dass das Q-Residuum des zu untersuchenden Spektrums zu den Q-Residuen des PCA-Modells passt. Ist die Nullhypothese wahr, kann eine annähernd normalverteilte Q-Residuum-Teststatistik erstellt werden.
- b. Das einstellbare Signifikanzniveau steuert die Wahrscheinlichkeit, mit der die Nullhypothese zurückgewiesen wird, falls sie wahr ist (bekannt als Typ-I-Fehler). Der Standardwert ist 5 %.
- c. Auf Basis dieser Teststatistik und des Signifikanzniveaus wird ein kritischer Wert für Q berechnet.
- d. Falls der Wert des Q-Residuums für das Spektrum grösser als der kritische Wert ist, wird die Nullhypothese zurückgewiesen und das Spektrum als potenzieller Ausreisser gekennzeichnet.

Ausreisserbewertung bei der Prädiktion (Quantifizierung)

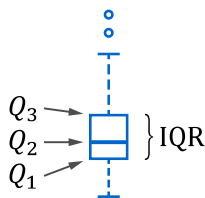
Bei der Quantifizierung steht während der Prädiktion eine Ausreisserbewertung zur Verfügung:

1. Vorbereitende Schritte:
 - a. Das Quantifizierungsmodell verwendet denselben Algorithmus wie oben. Die Basis ist jedoch das Quantifizierungsmodell (PLS-Modell) mit allen mittelwertzentrierten Spektren des Kalibrierdatensatzes. Datenvorbehandlung, Wellenlängenbereiche und Anzahl der latenten Variablen werden wie im Quantifizierungsmodell angegeben berücksichtigt.
 - b. Auf Basis des festgelegten Signifikanzniveaus berechnet das Quantifizierungsmodell die kritischen Werte für T^2 und Q (gestrichelte Linien im PLS-Influence-Plot). Die kritischen Werte werden im Quantifizierungsmodell hinterlegt.
2. Bei der Prädiktion werden die T^2 - und Q -Werte des Spektrums auf Basis des Quantifizierungsmodells berechnet.
3. Ist der T^2 - oder der Q -Wert des Spektrums grösser als der jeweilige kritische Wert, wird die Probe in Bezug auf das angewendete Quantifizierungsmodell als potenzieller Ausreisser betrachtet.

6.6 Referenzwert-Ausreisser – Algorithmus

Box-Plots ermöglichen die Erkennung von Ausreißern bei den Referenzwerten (*siehe "Referenzwert-Ausreißer (Quantifizierung)", Kapitel 4.3.4, Seite 59*).

Um die Schiefe der Verteilung zu berücksichtigen, werden die Ausreißergrenzwerte mit den folgenden Berechnungen angepasst.



Das **Medcouple** (MC) misst die Schiefe der Referenzwerte. Die Berechnung beginnt mit dem Median des Box-Plots, Q_2 . Mit allen möglichen Paaren (englisch *couples*) aus der oberen Hälfte und der unteren Hälfte der Referenzwerte wird eine Funktion berechnet. Der Median der Resultate ist das Medcouple:

$$\text{MC} = \text{med}_{y_i \leq Q_2 \leq y_j} \frac{(y_j - Q_2) - (Q_2 - y_i)}{y_j - y_i}$$

Hierbei entspricht Q_2 dem zweiten Quartil, das die Mittellinie im Box-Plot definiert, und y_i, y_i einem Paar von Referenzwerten.

Das Medcouple liegt immer zwischen -1 und 1 . Für eine symmetrische Verteilung ist $MC = 0$. Eine schiefe Verteilung mit $MC > 0$ ist zu den höheren Referenzwerten hin verzerrt, mit $MC < 0$ zu den niedrigeren Referenzwerten.

Die Berechnung der **angepassten Ausreissergrenzwerte** hängt davon ab, zu welcher Seite die Verteilung verschoben ist:

$$\text{MC} \geq 0: [Q_1 - 1.5 e^{-4\text{MC}} \text{IQR}; Q_3 + 1.5 e^{3\text{MC}} \text{IQR}]$$

$$\text{MC} < 0: [Q_1 - 1.5 e^{-3\text{MC}} \text{IQR}; Q_3 + 1.5 e^{4\text{MC}} \text{IQR}]$$

Bei einer symmetrischen Verteilung ($MC = 0$) betragen die Distanzen zwischen den Grenzwerten und der Box 1.5 IQR.

Die Exponentialfunktion ermöglicht eine genaue und robuste Ausreissererkennung bei verschiedenen Verteilungen mit unterschiedlicher Schiefe, wie empirisch nachgewiesen von M. Hubert und E. Vandervieren, *An adjusted boxplot for skewed distributions*, Computational Statistics & Data Analysis Band 52, Nr. 12 (Aug. 2008), S. 5186–5201.

Der erwartete Prozentsatz der markierten Ausreisser liegt bei etwa 1 % und ist dem Prozentsatz des Standard-Box-Plots für die Normalverteilung recht ähnlich. Hinweis: Dieser Prozentsatz ist unabhängig vom Signifikanzniveau.