



# statistiek

## Statistiek in het laboratorium van de ziekenhuisapotheek; deel 1.

M.C. de Brouwer<sup>1</sup>  
M.C.J. Langen<sup>1,2</sup>

<sup>1</sup> Laboratorium van de ziekenhuisapotheek Midden-Brabant  
Maria ziekenhuis  
Dr. Deelenlaan 5  
5042 AD Tilburg

<sup>2</sup> Correspondentie

### Samenvatting

Dit artikel behandelt de meest gebruikte statistische methoden in

laboratoria en hun toepassingen.

Aan bod komen toevallige en systematische fouten in analyses en de

wijze waarop deze kunnen worden geschat. Om een goede schatting van de fouten te kunnen geven worden eerst een aantal statistische grootheden beschreven, zoals de standaard deviatie en de variantie. Deze grootheden geven tevens een indicatie voor de grootte van de toevallige fouten. Voor het schatten van systematische fouten wordt gebruik gemaakt van een t-toets, waarbij twee verschillende waarden worden vergeleken. Voor het vergelijken van



toevallige fouten wordt de F-toets beschreven. De t-toets en de F-toets zijn voorbeelden van significantie toetsen.

Om resultaten over een concentratie gebied statistisch te toetsen wordt een regressie analyse beschreven, de methode van de kleinste kwadra-ten.

Hierbij wordt naast het berekenen van de regressielijn aandacht besteed aan fouten in de regressie-lijn en in de te bepalen monsters. Bij de fouten in de regressielijn komt tevens het beoordelen van de line-arieteit van de ijklijn door middel van de correlatie coëfficiënt en ANOVA-analyses aan de orde.

## Inleiding

Binnen laboratoria wordt het gebruik van statistiek steeds belang-rijker. Hierbij hoeft alleen maar te worden gedacht aan kwaliteitscon-troles. Ook het gebruik van de hoge-re statistiek, zoals in de chemome-trie, zal de komende jaren gemeen-goed gaan worden. De NVK-FAZ heeft dit onderkend door het organi-seren van statistiek cursussen voor zijn leden.

In twee artikelen zullen de basis principes van de statistiek (deel 1) worden uitgelegd en hoe deze wor-den toegepast binnen het laborato-rium van de ziekenhuisapotheek Midden-Brabant (deel 2).

Statistiek wordt binnen laboratoria gebruikt om de kwaliteit van resul-taten te toetsen.

Deze zijn afkomstig van onder ande-re controlemonsters [1], patiënt monsters en methode validaties [2]. Ieder soort analyse heeft zijn eigen criteria waarop getoetst wordt. Bij controlemonsters wordt gekeken of deze constant blijven over langere perioden, terwijl bij patiënt mon-sters het betrouwbaarheidsinterval rond de gevonden waarde van belang is.

In dit eerste deel wordt behandeld met welke statistische methoden resultaten kunnen worden getoetst en de uitkomsten geïnterpreteerd. Als leidraad voor dit artikel zijn de leerdoelen uit de statistiek cursus

van de NVK-FAZ en het bijbehoren-de cursusboek [3] gebruikt.

## Fouten in de analyse

Met behulp van statistiek worden systematische en toevallige fouten opgespoord en de grootte van deze fouten geschat.

Systematische fouten zijn er de oor-zaak van dat alle resultaten in dezelf-de mate afwijken van de werkelijke waarde. Het gemiddelde van een reeks metingen is hierbij niet gelijk aan de werkelijke waarde.

Systematische fouten zijn van invloed op de juistheid van een ana-lyse.

Toevallige fouten zijn er de oorzaak van dat de resultaten zowel boven als onder de werkelijke waarde lig-gen.

Er treedt een spreiding op in een reeks metingen. Toevallige fouten zijn van invloed op de precisie van de methode.

Een goede methode is precies en heeft een goede juistheid. Dit bete-kent dat toevallige en systematische fouten klein zijn.

In voorbeeld 1 worden aan de hand van een titratie de begrippen toege-licht.

In de praktijk kan het ook voorko-men dat moet worden gekozen tus-sen twee methoden, waarvan de ene precies is met een slechte juistheid (B) en een andere methode met een slechte precisie maar die wel juist is (C).

De methode met de slechte precisie kan een gemiddelde waarde geven die afwijkt van de werkelijke waarde. De methode met een goede precisie kan worden omgezet naar een methode met een goede juistheid als de systematische fout gevonden kan worden.

In de meeste gevallen zal de voor-keur dus uitgaan naar de methode met een goede precisie.

De keuze is echter ook afhankelijk van de snelheid, de kosten en ande-re factoren die bij een statistische analyse niet worden meegenomen.

Toevallige fouten worden geschat door statistische analyse van een reeks herhaalde metingen.

Systematische fouten kunnen alleen

## Voorbeeld 1

Vier analisten voeren een titratie uit op een 0,9% NaCl- oplossing. Het is bekend dat 10,00 mL titrant moet worden toegevoegd. In de tabel zijn de resultaten van de analis-ten gegeven.

| Analist             | Resultaten (ml) | Opmerkingen       |
|---------------------|-----------------|-------------------|
| A                   | 10,01           |                   |
|                     | 10,02           | goede precisie    |
|                     | 9,98            | goede juistheid   |
|                     | 9,99            |                   |
|                     | 9,99            |                   |
| B                   | 9,92            |                   |
|                     | 9,88            | goede precisie    |
|                     | 9,90            | slechte juistheid |
|                     | 9,91            |                   |
|                     | 9,89            |                   |
| C                   | 10,21           | slechte precisie  |
|                     | 10,11           | goede juistheid   |
|                     | 9,93            |                   |
|                     | 10,24           |                   |
|                     | 9,97            | slechte precisie  |
| D                   | 10,27           | slechte juistheid |
|                     | 10,01           |                   |
| Resultaten titratie |                 |                   |

Uit de tabel is af te leiden dat analist D de slechtste resultaten heeft en analist A de beste resultaten. Bij de analisten B en C ligt dit moeilijker.

op deze wijze bepaald worden als de werkelijke waarde exact bekend is. Dit is vaak niet het geval. Daarom wordt voor het bepalen van systematische fouten gebruik gemaakt van verschil toetsen. Hierbij worden de resultaten van een methode vergeleken met die van een andere onafhankelijke methode [2, 3].

## Statistische analyse van herhaalde metingen

Om de spreiding in een reeks metingen te bepalen wordt eerst het (rekenkundig) gemiddelde berekend. Dit is de som van alle metingen gedeeld door het totaal aantal metingen:

$$\bar{x} = \frac{1}{n} \cdot \sum_{i=1}^n x_i \quad (1)$$

$\bar{x}$  = het gemiddelde van alle metingen

$\sum_{i=1}^n x_i$  = de som van alle metingen

$n$  = aantal metingen

De spreiding in de metingen wordt vervolgens bepaald aan de hand van de standaard deviatie:

$$s = \sqrt{\frac{\sum (x_i - \bar{x})^2}{n-1}} \quad (2)$$

$s$  = standaard deviatie

De standaard deviatie geeft de spreiding als absolute grootheid. Om de precisie te kunnen vergelijken met resultaten met een andere grootheid kan de relatieve standaard deviatie (RSD) of variatie coëfficiënt worden gebruikt:

$$VC = \frac{s}{\bar{x}} \cdot 100\% \quad (3)$$

$VC$  = variatie coëfficiënt

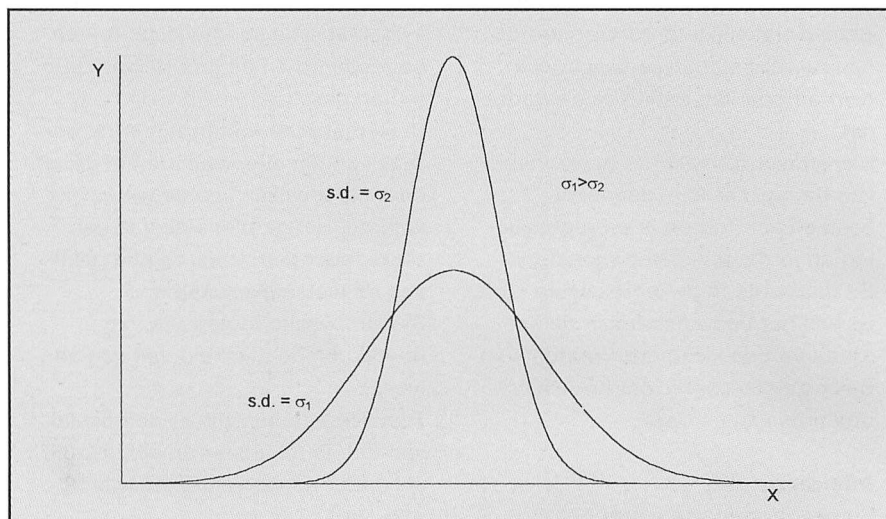
Een andere veel gebruikte statistische grootheid is de variantie. Dit is het kwadraat van de standaard deviatie:

$$v = s^2 \quad (4)$$

$v$  = variantie

In voorbeeld 2 wordt toegelicht hoe deze waarden in de praktijk worden toegepast. Als er een oneindig aantal metingen wordt gedaan ontstaat

tie. Het gemiddelde  $\bar{x}$  geeft een benadering van  $\mu$ . De spreiding rond  $\mu$  wordt gegeven door de standaard deviatie  $\sigma$  van de populatie. De



Figuur 1: Normaal distributie rond hetzelfde gemiddelde met verschillende standaard deviatie

een symmetrische curve rond  $\bar{x}$ . De verzameling van deze metingen wordt een populatie genoemd. Is er geen sprake van een systematische fout dan wordt het gemiddelde aangeduid met  $\mu$ , de werkelijke waarde voor het gemiddelde van de popula-

waarde van standaard deviatie,  $s$ , geeft een benadering van  $\sigma$ . Bij een grotere waarde van  $\sigma$  wordt de spreiding rond  $\mu$  groter, zoals weergegeven in figuur 1. Het model dat wordt gebruikt om een curve met een gemiddelde,  $\mu$ ,

### Voorbeeld 2

Twee analisten hebben een methode ontwikkeld voor de bepaling van lood in volbloed met behulp van de AAS. Om te kijken welke methode de beste precisie heeft voeren beide analisten zes maal een analyse uit op hetzelfde monster. Zij vinden de volgende resultaten.

|           |                                                                                                                                                         |
|-----------|---------------------------------------------------------------------------------------------------------------------------------------------------------|
| Analist A | 1,1 $\mu\text{mol/l}$ , 1,2 $\mu\text{mol/l}$ , 1,6 $\mu\text{mol/l}$ , 0,9 $\mu\text{mol/l}$ , 1,4 $\mu\text{mol/l}$ en 1,2 $\mu\text{mol/l}$ Pb       |
| Analist B | 248,64 $\mu\text{g/l}$ , 227,29 $\mu\text{g/l}$ , 248,64 $\mu\text{g/l}$ , 290,80 $\mu\text{g/l}$ , 331,52 $\mu\text{g/l}$ en 186,48 $\mu\text{g/l}$ Pb |

De standaard deviaties van deze beide experimenten geven:

|           |                          |                             |
|-----------|--------------------------|-----------------------------|
| Analist A | = 1,23 $\mu\text{mol/l}$ | $s = 0,242 \mu\text{mol/l}$ |
| Analist B | = 255,56 $\mu\text{g/l}$ | $s = 50,36 \mu\text{g/l}$   |

Als alleen naar de standaard deviaties wordt gekeken lijkt de methode van analist A een grotere precisie te hebben dan die van analist B. Analist A werkt echter met een andere grootheid dan B en het is daarom correcter om naar de variatiecoëfficiënt ( $vc$ ) te kijken.

|           |                                             |
|-----------|---------------------------------------------|
| Analist A | $vc = (0,242/1,23) \cdot 100\% = 19,6 \%$   |
| Analist B | $vc = (50,36/255,56) \cdot 100\% = 19,7 \%$ |

De  $VC$  voor beide methoden is vrijwel gelijk. Dit betekent dat de precisie van beide methoden ongeveer gelijk is. Bij het vergelijken van de precisie kan de standaard deviatie alleen maar gebruikt worden als dezelfde eenheden worden gebruikt en de uitkomsten dezelfde orde van grootte hebben.

en een standaard deviatie,  $\sigma$ , te beschrijven wordt normaal- of Gauss-distributie genoemd. Bij een normaal distributie blijkt dat 68% van de populatie ligt binnen een gebied van  $\pm 1\sigma$ , 95% binnen  $\pm 2\sigma$  en 99% binnen  $\pm 3\sigma$ . In de praktijk worden deze grenzen vaak toegepast bij het beoordelen van controle monsters [1, 3].

## Significantie toetsen

Significantie toetsen worden gebruikt om aan te tonen of een verschil tussen twee waarden significant is of dat het verschil veroorzaakt kan worden door toevallige fouten.

## Het vergelijken van gemiddelden

Bij een significantie toets wordt getoetst of een gemaakte hypothese juist is. Dit is de nulhypothese ( $H_0$ ). Bij het toetsen van de juistheid is de nulhypothese dat er geen systematische fout is. Het gemiddelde,  $\bar{x}$ , is gelijk aan de werkelijke waarde,  $\mu$ . De kans dat het gevonden gemiddelde gelijk is aan de theoretische waarde is klein. Bij een significantie toets wordt dan ook beoordeeld, hoe groot de waarschijnlijkheid is dat de afwijking afkomstig is van toevallige fouten. Gewoonlijk wordt gesteld dat de waarschijnlijkheid kleiner moet zijn dan 5%. Dit is het significantie niveau. Dit wordt genoemd als  $P = 0,05$ . Indien er een gro-

tere betrouwbaarheid gewenst is kan een hoger significantie niveau worden gebruikt, zoals  $P = 0,01$  of  $P = 0,001$ .

Om te bepalen of het verschil tussen  $\mu$  en  $\bar{x}$  significant is wordt student's  $t$  gebruikt.

$$t = \frac{|\bar{x} - \mu|}{\sqrt{n}/s}, \text{ waarbij } t \text{ altijd positief is.} \quad (5)$$

$t$  = student's  $t$

$$\sqrt{n}/s = \text{standaard deviatie in het gemiddelde} \quad (5b)$$

| Aantal vrijheidsgraden | Significantie niveau |        |        |
|------------------------|----------------------|--------|--------|
|                        | 0.1                  | 0.05   | 0.01   |
| 1                      | 6.314                | 12.706 | 63.657 |
| 2                      | 1.886                | 4.303  | 9.925  |
| 3                      | 1.638                | 3.182  | 5.841  |
| 4                      | 1.533                | 2.776  | 4.604  |
| 5                      | 1.476                | 2.571  | 4.032  |
| 6                      | 1.440                | 2.447  | 3.707  |
| 7                      | 1.415                | 2.365  | 3.499  |
| 8                      | 1.397                | 2.306  | 3.355  |
| 9                      | 1.383                | 2.262  | 3.250  |
| 10                     | 1.812                | 2.228  | 3.169  |

Tabel 1:  $t$ -tabel bij verschillende betrouwbaarheidsniveau's

De berekende waarde  $t$  wordt vergeleken met een kritische waarde voor  $t$ . Als de berekende waarde voor  $t$  groter is dan de kritische waarde wordt de nulhypothese verworpen. De kritische waarde voor een

bepaald significantie niveau en aantal vrijheidsgraden wordt opgezocht in tabel 1.

Het aantal vrijheidsgraden ( $v$ ) wordt bepaald door het aantal metingen ( $n$ ). Bij dit soort toetsen is het aantal vrijheidsgraden altijd  $n-1$ . In voorbeeld 3 wordt deze toets nader toegelicht. Uit vergelijking 5 kan tevens een betrouwbaarheidsinterval voor  $x$  worden afgeleid:

$$\text{betrouwbaarheidsinterval} = \bar{x} \pm \frac{t \cdot s}{\sqrt{n}} \quad (6)$$

Een andere wijze om de resultaten van een nieuwe methode te toetsen is door middel van een referentie methode. Hierbij worden twee gemiddelden ( $\bar{x}_1$  en  $\bar{x}_2$ ) met elkaar vergeleken. De nulhypothese is  $\bar{x}_1 - \bar{x}_2 = 0$

In het geval dat de standaard deviaties van de twee methoden niet significant verschillen wordt vergelijking 5 herschreven:

$$t = \frac{|\bar{x}_1 - \bar{x}_2|}{s / \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \quad (7)$$

waarbij  $s$  (de gepoolde standaard deviatie) wordt geschat met:

$$s = \sqrt{\frac{(n_1 - 1)s_1^2 + (n_2 - 1)s_2^2}{(n_1 - 1) + (n_2 - 1)}} \quad (8)$$

waarbij  $t$  ( $n_1 + n_2 - 2$ ) vrijheidsgraden heeft.

Indien de standaard deviaties niet gelijk zijn mogen vergelijking 7 en 8 niet gebruikt worden. Als in dit geval een verschil tussen de twee gemiddelden wordt gevonden kan dit ook worden veroorzaakt door het verschil in standaard deviaties.

Bij de beschreven methoden wordt het verschil tussen de gemiddelden getoetst in elke richting. Dit wil zeggen dat ervan uit wordt gegaan dat de te toetsen experimentele waarde zowel naar boven als naar beneden kan afwijken. Dit heet een tweezijdige toets.

In sommige gevallen wil men echter

## Voorbeeld 3

Bij de bepaling van lithium in serum met AES worden bij een referentie standaard die 1,00 mmol/l Li bevat de volgende waarden gemeten:

0,95, 0,96, 0,95, 0,97 mmol/l Li

Kan hier sprake zijn van een systematische fout?

Het gemiddelde van de waarden is 0,9575 mmol/l met een standaard deviatie van  $9,574 \cdot 10^{-3}$  mmol/l. Bij de nulhypothese dat er geen systematische fout optreedt en met gebruik van vergelijking 5 wordt gevonden:

$$t = \frac{|0,9575 - 1,00|}{\sqrt{4/9,574 \cdot 10^{-3}}} = 8,88$$

Uit tabel 1 kan worden afgelezen dat bij (4-1) vrijheidsgraden en  $P=0,05$  de kritische waarde voor  $t$  is 3,182. Er is hier met 95% betrouwbaarheid sprake van een systematische fout.

weten of de gevonden waarde alleen groter of alleen kleiner is. Dit is een één-zijdige toets. Bij een gegeven aantal vrijheidsgraden  $v$  en

| $v_1 \backslash v_2$ | 1    | 2    | 3    | 4    | 5    |
|----------------------|------|------|------|------|------|
| 1                    | 647  | 799  | 864  | 900  | 922  |
| 2                    | 38.5 | 39.0 | 39.2 | 39.3 | 39.3 |
| 3                    | 17.4 | 16.0 | 15.4 | 15.1 | 14.9 |
| 4                    | 12.2 | 10.6 | 9.9  | 9.6  | 9.3  |
| 5                    | 10.0 | 8.4  | 7.7  | 7.3  | 7.1  |

Tabel 2: F-tabel voor  $P = 0,05$

een bepaald betrouwbaarheidsniveau zal de kritische waarde bij een één-zijdige toets verschillen van een twee-zijdige toets. In dit geval wordt alleen gekeken naar de rechterkant van de Gauss-curve, als de experimentele waarde groter moet zijn, of naar de linkerkant van de Gauss-

## Het vergelijken van standaard deviaties

De t-toets wordt alleen toegepast voor het vergelijken van gemiddelden en dus voor het ontdekken van systematische fouten. In het geval dat de toevallige fouten van twee methoden, d.w.z. de standaard deviaties, worden getoetst moet een andere significantie toets worden toegepast, de F-toets.

Bij de F-toets wordt de verhouding tussen de varianties van twee reeksen metingen bepaald:

$$F_{v_1, v_2} = s_1^2 / s_2^2 \quad (9)$$

$v_1$  = aantal vrijheidsgraden ( $n_1 - 1$ )  
standaard deviatie teller

$v_2$  = aantal vrijheidsgraden ( $n_2 - 1$ )  
standaard deviatie noemer

en  $s_1^2$  en  $s_2^2$  worden zo gekozen dat  $F$  altijd  $\geq 1$  is. De nulhypothese is dat de varianties gelijk zijn. In dit geval zal de waarde van  $F$  rond de 1 liggen. Als de gevonden waarde voor  $F$

van het aantal metingen van beide experimenten, het significantie niveau en het soort toets. Dit is uitgewerkt in voorbeeld 4 [3].

## Regressie en Correlatie

De eerder beschreven statistische methoden zijn alleen van toepassing op een bepaald concentratie niveau. Als over een concentratie gebied wordt gemeten dan zijn deze methoden niet meer bruikbaar. Om de betrouwbaarheid van deze resultaten te toetsen wordt de regressie analyse toegepast.

Bij een regressie analyse wordt getoetst of er een verband bestaat tussen twee getallen reeksen. Dit verband kan bijvoorbeeld lineair, parabolisch of hyperbolisch zijn. In de analytische chemie wordt de regressie analyse vaak gebruikt om de lineariteit van een kalibratiecurve te bepalen, d.w.z. bestaat er een lineair verband tussen de concentratie van een bepaalde stof en de respons.

Indien er een lineair verband bestaat voldoet de gevonden kalibratiecurve aan:

$$y = bx + a \quad (10)$$

$b$  = de helling van de lijn  
 $a$  = het afsnijpunt van de lijn op de y-as

Veelal wordt de produkt-moment correlatie coëfficiënt gebruikt om de lineariteit van de kalibratiecurve te toetsen. Dit is echter niet terecht. De correlatie coëfficiënt geeft alleen aan of er een verband bestaat tussen de x- en y-waarden. Het zegt niets over het soort verband tussen de x- en y-waarden. De waarde van de correlatie coëfficiënt wordt bepaald volgens:

$$r = \frac{\sum \{(x_i - \bar{x})(y_i - \bar{y})\}}{\sqrt{\sum (x_i - \bar{x})^2 \cdot \sum (y_i - \bar{y})^2}} = \frac{S_{xy}}{\sqrt{S_{xx} \cdot S_{yy}}} \quad (11)$$

$r$  = correlatie coëfficiënt  
 $x_i$  = de concentratie van de component  
 $y_i$  = de respons bij concentratie  $x$   
 $\bar{x}$  = het gemiddelde van alle geme-

## Voorbeeld 4

Een methode is voorgesteld voor de bepaling van dexamethason in capsules. De nieuwe methode wordt vergeleken met de oude methode. De volgende resultaten werden verkregen:

|                | gemiddelde (mg) | Standaarddeviatie (mg) |
|----------------|-----------------|------------------------|
| oude methode   | 5,02            | 0,04                   |
| nieuwe methode | 5,02            | 0,02                   |

Voor elke bepaling werden 6 analyses uitgevoerd.

Om te bepalen of de precisie van de twee methoden verschillend is moet worden gekeken of de varianties significant verschillen. Hiertoe wordt  $F$  bepaald.

$$F_{5,5} = 0,04^2 / 0,02^2 = 4$$

Het aantal vrijheidsgraden voor iedere standaard deviatie wordt bepaald ( $n-1$ ). Aan de hand hiervan wordt in tabel 2 de kritische waarde voor  $F$  opgezocht; 7,15. Deze waarde is groter dan 4 en dus zijn de standaard deviaties, c.q. precisies, niet significant verschillend bij een betrouwbaarheidsniveau van 95%.

curve, als de experimentele waarde kleiner moet zijn. Daar een Gauss-curve symmetrisch is, is de waarschijnlijkheid de helft van die bij een twee-zijdige toets. De kritische waarde voor  $t$  in een één-zijdige toets ( $P=0,05$ ) wordt dan ook opgezocht onder  $P=0,10$  in tabel 1 voor een twee-zijdige toets [2,3].

groter is dan de kritische waarde voor  $F$  in tabel 2, dan wordt de nulhypothese verworpen. De F-toets kent, evenals de t-toets, een éénzijdige toets, methode A is preciezer dan methode B, en een tweezijdige toets, methode A en methode B hebben een verschillende precisie. De kritische waarde van  $F$  hangt af

## Voorbeeld 5

Bij de bepaling van ramipril met HPLC wordt ook het bijproduct diketopiperazine bepaald. De detectie vindt plaats met behulp van een UV/VIS-detector. Er wordt een kalibratiecurve opgenomen van 0,05 µg/mL tot 0,5 µg/mL. De gemeten resultaten zijn:

| Conc (µg/mL) | Respons (piekoppervlak) |
|--------------|-------------------------|
| 0,0511       | 2886                    |
| 0,1021       | 4861                    |
| 0,2042       | 6456                    |
| 0,3064       | 13200                   |
| 0,4085       | 17857                   |
| 0,5106       | 21539                   |

In de meeste gevallen zal de regressielijn worden berekend met behulp van een computer of calculator. Als de regressielijn manueel wordt berekend, wordt eerst een tabel gemaakt met de benodigde gegevens.

| $x_i$             | $y_i$ | $x_i - \bar{x}$ | $(x_i - \bar{x})^2$ | $y_i - \bar{y}$ | $(y_i - \bar{y})^2$ | $(x_i - \bar{x})(y_i - \bar{y})$ |
|-------------------|-------|-----------------|---------------------|-----------------|---------------------|----------------------------------|
| 0,0511            | 2886  | -0,2127         | 0,0452              | -8747           | $7,651 \cdot 10^7$  | 1860                             |
| 0,1021            | 4861  | -0,1617         | 0,0261              | -6772           | $4,586 \cdot 10^7$  | 1095                             |
| 0,2042            | 9456  | -0,0596         | 0,0036              | -2177           | $0,474 \cdot 10^7$  | 130                              |
| 0,3064            | 13200 | 0,0426          | 0,0018              | 1567            | $0,246 \cdot 10^7$  | 67                               |
| 0,4085            | 17857 | 0,1447          | 0,0209              | 6224            | $3,874 \cdot 10^7$  | 900                              |
| $\Sigma$ : 1,5829 | 69799 | 0               | 0,1585              | 0               | $26,64 \cdot 10^7$  | 6497                             |

$$\bar{x} = 1,5829/6 = 0,2638; \bar{y} = 69799/6 = 11633$$

Als eerste wordt de correlatie coëfficiënt berekend:

$$r = 6497 / \sqrt{0,1585 \cdot 26,64 \cdot 10^7} = 0,9995$$

De waarde van r ligt dichtbij 1 en er kan van worden uitgegaan dat er een positieve correlatie is tussen x en y.

Vervolgens wordt de vergelijking van de lijn berekend:

$$b = 6497 / 0,1585 = 40967$$

$$a = 11633 - (40967 \cdot 0,2638) = 826$$

De vergelijking van de kalibratiecurve is:

$$y = 826 + 40967x, r = 0,9995$$

$$b = \frac{\sum \{(x_i - \bar{x})(y_i - \bar{y})\}}{\sum (x_i - \bar{x})^2} = \frac{S_{xy}}{S_{xx}} \quad (12)$$

en

$$a = \bar{y} - b\bar{x} \quad (13)$$

Deze lijn is de regressielijn van y over x en geeft aan hoe y varieert ten opzicht van x. De regressielijn van x over y is niet dezelfde lijn. Die lijn gaat er van uit dat alle fouten optreden in de x-waarden. In voorbeeld 5 wordt een uitwerking gegeven.

De regressielijn wordt vaak gebruikt om de concentratie van een monster te bepalen. De toevallige fouten in de waarden voor de helling en het afsnijpunt zijn dus belangrijk om een goede indicatie van de betrouwbaarheid van de resultaten te krijgen. Om de fout in de helling en het afsnijpunt te berekenen wordt eerst de standaard deviatie van de regressielijn berekend:

$$s_{\hat{y}} = \sqrt{\frac{\sum (y_i - \hat{y}_i)^2}{n-2}} \quad (14)$$

$s_{\hat{y}}$  = standaard deviatie regressielijn

$\hat{y}_i$  = berekende waarde voor y op de regressielijn

$y_i$  = gemeten waarde voor y

ten concentraties  
 $\bar{y}$  = het gemiddelde van alle gemeten y-waarden

De waarde van r kan tussen -1 en +1 liggen. Een r-waarde van -1 beschrijft een perfecte negatieve correlatie en een r-waarde van +1 beschrijft een perfecte positieve correlatie. Als de waarde van r dicht bij nul komt neemt de correlatie tussen de x- en y-waarden af.

Als wordt aangenomen dat er een correlatie bestaat tussen x en y dan kan het soort verband bepaald worden. Om de beste rechte lijn door de meetpunten te bepalen wordt veelal

de methode van de kleinste kwadraten toegepast. Hierbij wordt ervan uitgegaan dat er alleen een fout optreedt in de respons (y-waarden). Hieruit volgt dat voor de beste lijn door de punten het verschil tussen de gevonden y-waarde en de berekende lijn minimaal moet zijn. Omdat sommige punten een positief en andere een negatief verschil geven, wordt het kwadraat van het verschil genomen. De som van de kwadraten van de verschillen moet zo klein mogelijk zijn. Bij deze methode gaat de lijn altijd door het punt  $(\bar{x}, \bar{y})$ . De vergelijking van de lijn wordt berekend met:

Met behulp van  $s_{x/y}$  kunnen de standaard deviaties in de helling en in het afsnijpunt worden berekend

$$S_b = \frac{s_{x/y}}{\sqrt{\sum (x_i - \bar{x})^2}} \quad (15)$$

$S_b$  = standaard deviatie helling

en

$$S_a = s_{x/y} \cdot \sqrt{\frac{\sum x_i^2}{n \sum (x_i - \bar{x})^2}} \quad (16)$$

$S_a$  = standaard deviatie afsnijpunt  
De betrouwbaarheidsintervallen voor de helling en het afsnijpunt

# Voorbeeld 6

Van de kalibratie curve uit voorbeeld 5 wil men het betrouwbaarheidsinterval van de helling en van het afsnijpunt berekenen.

Evenals bij het berekenen van de vergelijking van de ijklijn wordt eerst een tabel gemaakt

| $x_i$      | $x_i^2$ | $y_i$  | $\hat{y}_i$ | $y_i - \hat{y}_i$ | $(y_i - \hat{y}_i)^2$ |        |
|------------|---------|--------|-------------|-------------------|-----------------------|--------|
| 0,0511     | 0,0026  | 2886   | 2918        | -32               | 1024                  |        |
| 0,1021     | 0,0104  | 4861   | 5010        | -149              | 22201                 |        |
| 0,2042     | 0,0416  | 6456   | 9193        | 266               | 70756                 |        |
| 0,3064     | 0,0938  | 13200  | 13377       | -177              | 31329                 |        |
| 0,4085     | 0,1669  | 17857  | 17560       | 297               | 88209                 |        |
| 0,5106     | 0,2516  | 21539  | 21744       | -205              | 42025                 |        |
| $\Sigma$ : | 1,5829  | 0,5253 | 66799       | 69802             | 0                     | 255544 |

Als eerste wordt de standaard deviatie van de kalibratie curve,  $S_{x/y}$ , uitgerekend (vgl. 14).

$$s_{x/y} = \sqrt{255544 / (6 - 2)} = 253$$

Vervolgens wordt de standaard deviatie in de helling en in het afsnijpunt bepaald (vgl. 15 en vgl. 16).

$$s_b = 253 / \sqrt{0,1585} = 635$$

$$s_a = 253 \sqrt{(0,5254 / (6 \cdot 0,1585))} = 197$$

Om het betrouwbaarheidsinterval te kunnen bereken wordt de t-waarde opgezocht in tabel 1 voor (n-2 = 4) vrijheidsgraden en P= 0,05. Deze waarde is 2,776.

$$b = 40967 \pm 2,766 \cdot 853 = 40967 \pm 1763$$

$$a = 826 \pm 2,766 \cdot 197 = 826 \pm 546$$

Het belangrijkste wat uit deze gegevens valt op te maken is dat de fout in de helling gering is en dat de lijn met een waarschijnlijkheid van 95% niet door 0 gaat. De ondergrens van het betrouwbaarheidsinterval van a is 826-548= 280.

waarde kleiner wordt naar mate de waarde voor  $y_0$  nadert tot  $\bar{y}$ . De algemene vorm van het betrouwbaarheidsinterval voor een berekende concentratie is schematisch weergegeven in figuur 2.

Als gebruik wordt gemaakt van de standaard additie methode wordt de regressielijn op dezelfde wijze berekend, maar voor het berekenen van de fout in x worden de formules 21 en 22 gebruikt.

$$s_{x_E} = \frac{s_{x/y}}{b} \sqrt{\frac{1}{n} + \frac{\bar{y}^2}{b^2 \sum (x_i - \bar{x})^2}} \quad (21)$$

$$\text{betrouwbaarheidsinterval} = x_E \pm t_{s_{x_E}}, (n-2) \text{ vrijheidsgraden} \quad (22)$$

$x_E$  = berekende waarde voor x

Zoals eerder vermeld is met behulp van de correlatie coëfficiënt snel te bepalen of er een correlatie bestaat tussen twee getal reeksen, maar zegt dit niets over het verband tussen de reeksen. Een methode om te testen of er een mogelijk lineair verband bestaat is door gebruik te maken van een variantie analyse (Analysis of Variance =ANOVA). ANOVA berekeningen vereisen veel rekenwerk en daarvoor wordt bij ANOVA-analyses gebruik gemaakt van computerprogramma's. Binnen het laboratorium van de ziekenhuis-apotheek Midden-Brabant wordt hiervoor gebruik gemaakt van een Excel® applicatie.

Het gaat in het kader van dit artikel te ver om uitgebreid in te gaan op de berekeningen die plaats vinden bij een ANOVA-analyse. Van belang zijn twee uitkomsten die worden verkregen, de Goodness of Fit (GOF) en de Lack of Fit (LOF). De GOF geeft aan hoe goed de ijklijn voldoet aan het lineaire model. Naarmate de waarde voor GOF groter wordt, is de waarschijnlijkheid groter dat er een lineair verband is. De waarde van GOF wordt bepaald met behulp van de F-toets, is deze waarde groter dan de waarde uit de tabel dan is er een grote waarschijnlijkheid dat er een lineair verband bestaat.

worden berekend volgens:

$$\text{betrouwbaarheidsinterval helling} = b \pm t s_b \quad (17)$$

$$\text{betrouwbaarheidsinterval afsnijpunt} = a \pm t s_a \quad (18)$$

waarbij de t-waarde voor het gewenste betrouwbaarheidsniveau en (n-2) vrijheidsgraden wordt genomen. Een rekenvoorbeeld van de bepaling van het betrouwbaarheidsinterval van helling en afsnijpunt is gegeven in voorbeeld 6.

Uit de berekende regressielijn kan op eenvoudige wijze een x-waarde worden berekend bij een gemeten y-waarde. Hiervoor worden zowel de helling als het afsnijpunt gebruikt. Deze waarden hebben echter een fout en als gevolg hiervan zit er ook

een fout in de berekende x-waarde. Het betrouwbaarheidsinterval van x wordt berekend met:

$$s_{x_0} = \frac{s_{x/y}}{b} \sqrt{\frac{1}{m} + \frac{1}{n} + \frac{(y_0 - \bar{y})^2}{b^2 \sum (x_i - \bar{x})^2}} \quad (19)$$

$$\text{betrouwbaarheidsinterval} = x_0 \pm t s_{x_0}, (n-2) \text{ vrijheidsgraden} \quad (20)$$

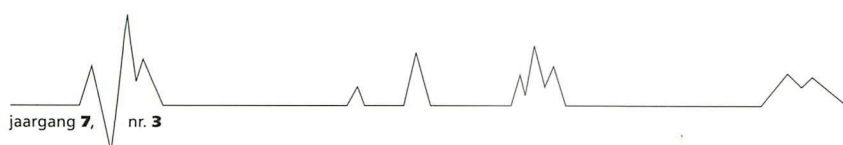
$x_0$  = berekende waarde voor x

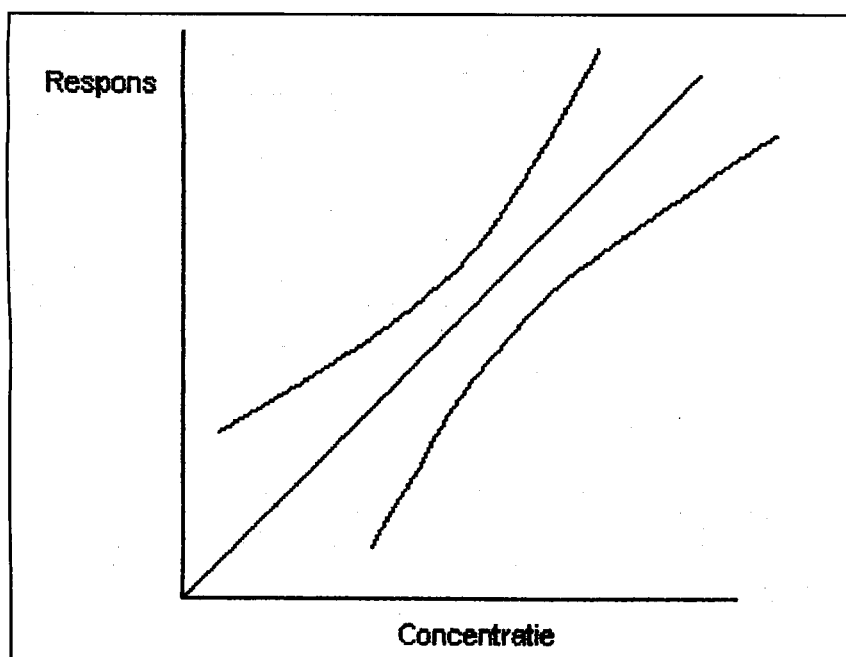
$y_0$  = gemeten y-waarde

m = aantal onafhankelijke metingen van  $y_0$

In het geval van een enkelvoudige meting geldt dat m=1, voor onafhankelijke duplo metingen geldt dat m=2, enz.

Uit vergelijking 19 volgt dat het betrouwbaarheidsinterval voor de x-





Figuur. 2: Algemene vorm voor het betrouwbaarheidsinterval rond een regressielijn

De LOF geeft een maat voor de afwijking van het lineaire model. Als de waarde voor LOF kleiner wordt, is de afwijking ten opzichte van het lineaire model kleiner. De waarde van LOF wordt eveneens bepaald met behulp van de F-toets, is deze

waarde kleiner dan de waarde uit de tabel dan vertoont het model geen lack of fit.

## Nawoord

Met statische methoden is een goed inzicht te krijgen in de betrouwbaarheid van gevonden resultaten. De

toepassing van statistiek binnen laboratoria is echter nog gering. Software op apparatuur biedt vaak wel de mogelijkheid om ijklijnen met een correlatie coëfficiënt uit te rekenen. Verdere statistiek wordt niet bedreven. Ook bij interpolatie van meetwaarden naar concentraties ontbreekt een schatting van de mogelijke fout in het analyse resultaat. Voor analisten is het echter van belang te weten wat voor waarde moet worden gehecht aan een gevonden resultaat. In het tweede deel worden een aantal praktijk voorbeelden behandeld, zoals deze zijn voorgekomen in het laboratorium van de ziekenhuis apotheek Midden-Brabant.

## Referenties

1. C.M.M. Haring, *Ziekenhuisfarmacie*, 6(1990), 83-86
2. F.J. van de Vaart, *Extract*, 2(1993), 13-23
3. Miller and Miller, *Statistics for Analytical chemists*, 3th edition, Ellis Horwood and Prentice Hall, 1994
4. Drs. A. Buijs, *Statistiek om mee te werken*, derde, herziene druk, Stenfert Kroese, 1989