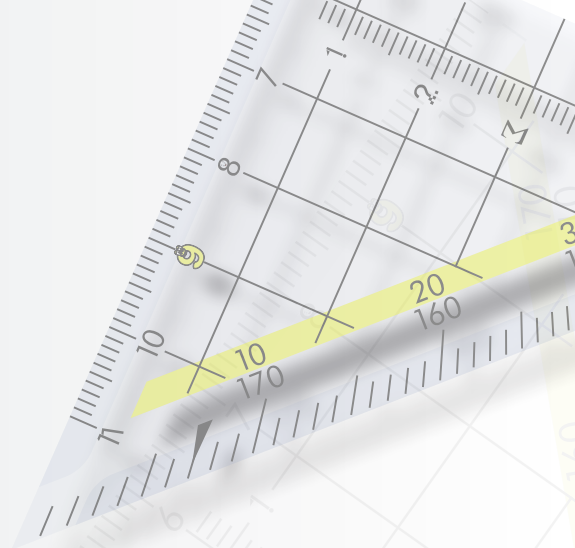
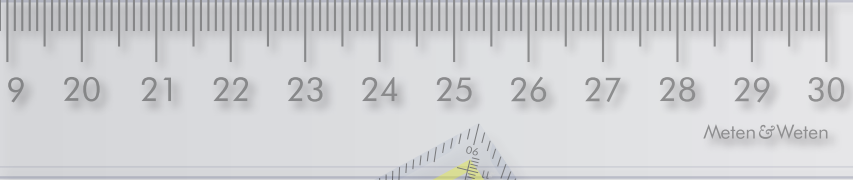
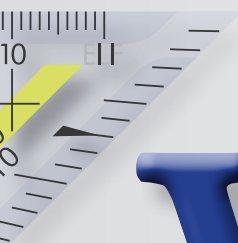




Meten & Weten



Gerard Tuk

Meten & Weten

Meten & Weten

Gerard Tuk

Meten & Weten

Uitgave 2024/1

Eerste uitgave: 27 december 2023

Tweede uitgave: 26 augustus 2024

Auteur: Gerard Tuk

Druk: Pumbo.nl

Omslagontwerp: Gerard Tuk

Vormgeving binnenwerk: Gerard Tuk

Lettertype: Linux Libertine (tekst)

Amaranth (koppen)

Consolas (code)

Alles uit deze uitgave mag worden verveelvoudigd, door middel van druk, fotokopieën, geautomatiseerde gegevensbestanden of op welke andere wijze ook, ook zonder voorafgaande schriftelijke toestemming van de uitgever.

Inhoud

Amuses	11
0. Lees wijzer	13
1. “Tweeënzeventig zes!”	17
2. Snelheidscontroles	20
3. De sleutel	23
4. Doellijntechnologie	27
5. Onzeker zijn	31
6. W’rom?	33
7. Model & Werkelijkheid	35
Fouten & Onzekerheden	37
8. Model van een meting	39
9. Beste schatting $\pm$ onzekerheid [eenheid]	43
10. Systematisch & Toevallig	45
11. Toevallige systematische fouten	49
12. Nauwkeurig & Precies	51
13. Nauwkeuriger & Preciezer	55
14. Spreiding in ruimte en tijd	57
15. Spreiding in wat we meten	59
16. Spreiding door het meten	61
17. Beïnvloeden van wat we meten	63
18. Schietschijven en rozen	64
19. Onzekerheidsrelatie van Heisenberg	66
20. Error & Uncertainty	67
21. Werkelijk & Geaccepteerd	69
22. Beste schatting van de waarde	71
23. Beste schatting van de precisie	73
24. Beste schatting (Pythonsimulatie)	75
25. Fout gebruik van de term ‘error’	77

Notatie	81
26. Afronden	83
27. Significante cijfers	86
28. Significante cijfers in berekeningen	88
29. Significante cijfers in onzekerheden	91
30. Noteren van onzekerheden.....	93
31. De Grote Negen.....	97
32. Verschillen meten	101
33. Digitaal & Analooq.....	103
34. Aflezen van een analoge schaal.....	106
35. Aflezen van een digitale schaal.....	109
Foutdoorwerking.....	111
36. Foutdoorwerking bij optellen.....	113
37. Doorwerking bij aftrekken	115
38. Een schuine tafel	117
39. Doorwerken bij vermenigvuldigen.....	119
40. En als je het niet mag verwaarlozen?.....	123
41. Doorwerken in het algemeen.....	125
42. Doorwerken bij delen	129
43. Doorwerken bij een exacte factor	131
44. Doorwerken bij machtsverheffen	132
45. Doorwerken bij Louis Lyons	133
Afhankelijk & Onafhankelijk.....	137
46. Voorbeelden van (on)afhankelijkheid.....	139
47. Extreme voorbeelden	142
48. “Een getal onder de tien?!”	145
49. Onafhankelijke fouten optellen.....	147
50. Met én zonder afhankelijkheid.....	149
51. (On)afhankelijke normale verdelingen.....	151

Begrippen.....	155
52. Onzekerheid	157
53. Verificatie & Validatie	160
54. Relatief & Absoluut	161
55. Werkelijke waarde.....	163
56. Fout.....	165
57. Type A & Type B	167
58. Meetschalen.....	169
59. Discrepantie	171
60. Relevant & Significant	174
61. Onzekerheid over de onzekerheid	176
62. Effectief & Efficiënt	177
63. Kwaliteit.....	179
64. Procenten & Procentpunten	181
65. Absolute & relatieve onzekerheden.....	183
66. Kleine waarde, grote onzekerheid.....	185
Statistiek.....	187
67. ‘The Average’	189
68. Gemiddelde(n).....	191
69. Standaardafwijking.....	193
70. Een klein aantal metingen	195
71. Eindig aantal mogelijkheden.....	199
72. Kansvariabele.....	201
73. Verwachtingswaarde.....	205
74. Permutaties & Combinaties.....	207
75. Eén dobbelsteen.....	211
76. Betere schatting van spreiding	215
77. De spreiding in de spreiding	218
78. Rough-and-ready approach.....	221

79. Kansverdeling resultaat en meting	225
80. Overschrijdingskans en significantie	230
81. Verwaarloosbare kans.....	233
Toevallige fouten	235
82. Systematische fout door toeval (1).....	237
83. Systematische fout door toeval (2).....	239
84. Systematisch toenemende fout	242
Chauvenet.....	245
85. Het criterium van Chauvenet.....	247
86. Chauvenet in andere boeken.....	251
87. Overschrijdingskansen	253
88. Overschrijdingskansen online	257
89. Chauvenet in Excel.....	258
90. Chauvenet in Python	260
Getallen	263
91. Zwevendekommagetallen	265
92. De IEEE 754 Converter	268
93. Decimalen in Python en C	270
94. Exacte getallen in Python	273
Grafieken.....	277
95. Lijnen door punten	279
96. Kleinste kwadraten.....	282
97. Résidu	285
98. Raad het volgende getal.....	288
99. Past het model?.....	289
Rekenen.....	293
100. De taylorreeks van een sinus	295
101. Graden	297
102. Radialen.....	299

103. 3x.....	301
104. Maak er maar geen punt van	303
Weten.....	307
105. Herhalen van een meting.....	309
106. Falsificatie	311
107. Inschatten wat verwaarloosbaar is	313
108. Vroeg vinden van verbeter(d)ingen.....	315
109. Proberen	317
110. Pas op.....	319
Antwoorden	321
Bijlagen.....	351
A. Alt-codes en Unicodes	353
B. Bekende afgeleiden.....	355
C. Competenties bij het meten.....	357
D. Doorwerking van fouten.....	359
E. Enige SI-voorvoegsels	362
F. Foutenvinders	363
G. Gebruikte boeken.....	364

Amuses

0. Lees wijzer

Toen ik in augustus 1986 werd gekeurd voor militaire dienst, vonden ze het bij Defensie nodig of nuttig om mijn lengte vast te stellen. Mijn naam werd afgeroepen en er bulderde een bevel: ik diende op blote voeten met mijn rug tegen de muur te gaan staan. Als dat maar goed afliep.

Zo begint straks hoofdstuk 1, en hoe het afloopt, zie je daarna wel.

Vóór hoofdstuk 1 komt hoofdstuk 0, met een leeswijzer. Dit boek is namelijk anders dan de meeste andere boeken. Het heeft een beetje de stijl van *Lessons Learned in Software Testing* van Cem Kaner. Die schrijft in zijn voorwoord:

Imagine that you are holding a bottle of 50-year-old port. There is a way to drink port. It is not the only way, but most folks who have enjoyed port for many years have found some guidelines that help them maximize their port-drinking experience. Here are just a few:

Lesson 1: Don't drink straight from the bottle. (...) Don't gulp down the port.

Lesson 2: Don't drink the entire bottle. If you are drinking because you are thirsty, put down the port and drink a big glass of water.

Lees wijzer! Dit boek is wel bedoeld om van voor tot achter door te lezen, maar niet in één adem. De opzet is dat je losse 'lesjes' leest en tot je door laat dringen. Niet uit de fles drinken, maar de tijd nemen om het te laten indalen. De oefeningen dienen daartoe: je gaat het beter snappen en onthouden als je zelf iets doet. Zwemmen en pianospelen leer je ook niet uit een boek.

Meten & Weten bestaat eigenlijk uit verhaaltjes en heeft hier en daar onbelangrijke uitstapjes. Die zijn grijs gemarkeerd, zoals deze alinea. Die mag je rustig overslaan, maar misschien zijn ze juist wel leuk. Of toch leerzaam, omdat ze je aanzetten tot verder nadenken over de stof.

Het hele onderdeel Amuses mag je ook overslaan.

Studieboeken gebruik je niet altijd voor je lol, maar het is mooi meegenomen als ze prettig zijn om te lezen. De schrijver van *Metten & Weten* heeft daar naar gestreefd: serieuze inhoud op een luchtige manier gebracht.

Bij de opleiding Mechatronica aan de Haagse Hogeschool hebben er voor het vak *Error Budgeting* twee verschillende boeken op de lijst gestaan, namelijk die van Taylor (1997) en van Hughes & Hase (2010). Hoe begonnen die? Wat waren hun eerste zinnen? Werd het meteen al vanaf het prille begin een thriller? En hoe ging het verder?

De ouverture van (het voorwoord van) het boek van Taylor luidt als volgt.

All measurements, however careful and scientific, are subject to some uncertainties.

De eerste alinea van (het eerste hoofdstuk van) het boek van Hughes & Hase opent met een vraag en een antwoord uit een bron met aanzien en status.

What is the role of experiments in the physical sciences? In his Lectures on Physics, Volume 1, page 1, the Nobel Laureate Richard Feynman states (Feynman 1963):

The principle of science, the definition, almost, is the following: The test of all knowledge is experiment. Experiment is the sole judge of scientific ‘truth’.

Daarna wordt er een belangrijke uitspraak gedaan.

We will emphasise in this book that an experiment is not complete until an analysis of the uncertainty in the final result to be reported has been conducted.

De vraag naar de rol van experimenten wordt gevolgd door een antwoord over natuurwetenschappelijke rechtvaardiging. En dan over wat er onlosmakelijk aan een experiment verbonden is: een analyse van de onzekerheid.

Berendsen (2016) zet in “zijn” hoofdstuk 1 in met een onmogelijkheid:

It is impossible to measure physical quantities without errors.

Ratcliffe & Ratcliffe (2014) steken op een soortgelijke manier van wal.

Every time anyone takes a measurement, the only certainty is that the measurement is not exact.

Dat klinkt nog wat cynischer, wanhopig bijna. Maar ook filosofisch. Het heeft iets van een milde spot in zich. Je hebt maar één zekerheid: het klopt niet helemaal. En dat is nog universeel ook. Het geldt altijd, voor eeuwig en voor iedereen.

Kirkup & Frenkel (2006) kiezen hun uitgangspunt in het dagelijkse leven en hebben het niet meteen over metingen.

We live with uncertainty every day. Will the weather be fine for a barbecue at the weekend? What is the risk to our health posed by a particular item of diet or environmental pollutant? Have we invested our money wisely? It is understandable that we would like to be able to eliminate, or at least reduce, uncertainty. If we can reduce it significantly, we become more confident that a desirable event will happen, or that an undesirable event will not. To this end we seek out accredited professionals, such as weather forecasters, medical researchers and financial advisers.

Is dit niet een heel ander soort metingen? De onzekerheid die ontstaat doordat je geen glazen bol hebt, is dat niet iets anders dan een meting doen in het hier en nu en nog niet precies weten wat je hebt? Tot op zekere hoogte wel. Maar het weer van morgen is niet zuiver toeval. De onzekerheden in de metingen van vandaag maken dat we niet alles weten over het volgende etmaal.

Onzekerheid is iets waar je last van hebt en wat ingeperkt moet worden. Je zou kunnen kiezen voor een hapje en een drankje binnenshuis in plaats van een barbecue. Je verstand gebruiken bij wat je eet en drinkt en accepteren dat ziekte bij het leven hoort. En je geen zorgen maken over beleggingen, maar blij zijn dat je geen schulden hebt.

Over schulden gesproken: Lyons (1991) begint zijn boek met

שְׂגִיאוֹת מִי-יָבִין

Mochten er lezers zijn die geen Hebreeuws beheersen – die heb je er altijd bij: dit is een Bijbeltekst. Psalm 19:13.

Maar wie kan al zijn fouten kennen?

Wéér over fouten dus. Maar niet het soort waar *Metén* & *Weten* over gaat.

Spreek mij vrij van verborgen zonden.

‘Zonden’ zijn ook fouten, maar wel van een heel andere categorie.

Dezelfde Louis Lyons schrijft op de achterkant van zijn boek het volgende.

It is usually straightforward to calculate the result of a practical experiment in the laboratory. Estimating the accuracy of that result is often regarded by students as an obscure and tedious routine, involving much arithmetic. An estimate of the error is, however, an integral part of the presentation of the results of experiments.

Schatten van nauwkeurigheden en/of onzekerheden is kennelijk nog niet eens leuk ook! De woorden *obscure* en *tedious* voorspellen niet veel goeds.

Taylor vervolgt zijn voorwoord wat dat betreft iets optimistischer.

Error analysis is the study and evaluation of these uncertainties, its two main functions being to allow the scientist to estimate how large his uncertainties are, and to help him to reduce them when necessary. The analysis of uncertainties, or “errors”, is a vital part of any scientific experiment, and error analysis is therefore an important part of any college course in experimental science. It can also be one of the most interesting parts of the course. The challenges of estimating uncertainties and of reducing them to a level that allows a proper conclusion

to be drawn can turn a dull and routine set of measurements into a truly interesting exercise.

‘De uitdagingen’ maken wat saai is een ‘waarlijk interessante oefening.’

Hughes & Hase (2010), de auteurs van het boek dat op de lijst stond vóór *Metten & Weten*, hebben het niet zozeer over ‘saai’ of ‘leuk’, maar geven al vroeg in hun eerste hoofdstuk aan dat het *belangrijk* is wat we doen. Je bent domweg niet klaar met je meting of experiment als je niets gezegd hebt over hoe (on)zeker je van je zaak bent.

We will emphasise in this book that an experiment is not complete until an analysis of the uncertainty in the final result to be reported has been conducted.

Important questions such as:

- *do the results agree with theory?*
 - *are the results reproducible?*
 - *has a new phenomenon or effect been observed?*
- can only be answered after appropriate error analysis.*

De vraag is niet eens of het leuk is. Het *moet*, want anders weet je zo weinig.

Mijn gewaardeerde collega Mark Schrauwen schreef me toen hij Editie 2023 van *Metten & Weten* aan het lezen was:

*Ik vind het lastig om te bepalen met wat voor boek ik te maken heb. Het ene moment denk ik het is een verhalend boek zoals *Zoekt & Gijzelt*, het andere moment is het iets meer een studieboek.*

Dat snapte ik wel. Sterker nog: *Metten & Weten* is ook een verhalend boek én een studieboek, maar voorzag dus niet volledig in zijn behoefte.

Ik zou zelf gebaat zijn met een soort van mindmap die een grafische relatie tussen de hoofdstukken aanbrengt. Die ga ik sowieso voor mezelf maken.

Wat een goed idee! Het boek nemen zoals het is, en er dan *zelf* mee aan de slag gaan. Zonder leeswijzer is het devies: Lees wijzer!

Oefeningen

1. Geef in maximaal vijftig woorden de essentie weer van wat diverse boeken in hun openingszinnen te melden hebben.
2. Bedenk wat je gaat doen om ervoor te zorgen dat je tijdens het lezen van dit boek kunt denken: “Ik lees wijzer”.

1. “Tweeënzeventig zes!”

Toen ik in augustus 1986 werd gekeurd voor militaire dienst, vonden ze het bij Defensie nodig of nuttig om mijn lengte vast te stellen. Mijn naam werd afgeroepen en er bulderde een bevel: ik diende op blote voeten met mijn rug tegen de muur te gaan staan. Als dat maar goed afliep.

Iemand liet een stompe staaf op mijn – toen nog – donkere haren dalen, las iets af bij een of andere streep, en schreeuwde: “tweeënzeventig zes!”

De toon en het bijpassende volume maakten duidelijk dat eventuele tegenspraak geen gevolgen zou hebben. Of juist wel.

Dat “tweeënzeventig zes!” duurde, als je de galm niet meerekende, ongeveer 1,414213562 seconde. De beide woorden riepen bij mij een paar gedachten op, die ik strak zwijgend voor me hield. De eerste was dat er een *eenheid* ontbrak, maar die sprak kennelijk vanzelf. Een belangrijkere bespiegeling was dat die *zes* waarschijnlijk het aantal millimeters betrof. Het leek me onwaarschijnlijk dat je zó goed kon vaststellen hoe lang iemand was. Wat me nog het meest verbaasde: dat ze er *precies* een meter naast zaten...

Om goedgekeurd te worden, moest je minimaal 1 meter 60 zijn en maximaal 2 meter. Was dat toegestane bereik van 40 cm soms twee keer tweemaal de standaardafwijking? In dat geval was $\sigma = 10$ cm en viel ik in de middengroep. Dat had ik niet verwacht, omdat ik vroeger bij gym altijd tot de minst lange leerlingen behoorde. Als die dertig jongens op volgorde van lengte op de witte lijn moesten gaan staan, zat ik altijd bij de kleinste vijf.

Die millimeters leken mij wat overdreven en die 1 en de eenheid lieten ze waarschijnlijk weg voor het gemak. Verdedigbare keuzes, van Defensie. Een afstand van top tot teen druk je nu eenmaal uit in meters en centimeters. Waarom zou je het erbij zeggen, als het vanzelf spreekt? Bij veruit de meeste zeventienjarigen is het cijfer voor de komma een 1. Zo’n twee procent van de mensen met mijn geboortejaar haalt de twee meter.

Als vwo’ertje aan het begin van zijn laatste schooljaar heb ik die keuring zó serieus genomen dat ik de afgeronde 1,73 m jarenlang als mijn officiële lengte heb beschouwd. Paspoorten, rijbewijzen, ID-kaarten: meer dan dertig jaar lang hebben die het resultaat van deze ene meting vermeld. Nog steeds trouwens. Op het gemeentehuis werd ik bij het verlengen van een officieel document nooit eens een keertje nagemeten. Ze vroegen gewoon of het oude nog klopte en ik zei “ja”. Zonder dat zelf gecontroleerd te hebben. Ik hield me met haast militaire discipline aan wat die officier als officieel had vastgelegd.

In de zomer van 2022 besloot ik toch maar eens een poging te doen zelf vast te stellen hoe ver ik boven het maaiveld uitkwam. In tweevoud, omdat herhalen van een meting vaak verstandig is. Beide keren kwam ik op 1,71 m... Hadden ze dan niet goed gekeken toen, in die militaire meetkamer in Breda? Of gaan mensen al veel eerder krimpen dan ik dacht?

Studenten van nu zijn gemiddeld ruim een centimeter langer dan “in mijn tijd”. (Bron: een oppervlakkig zoektochtje via Google. Doe dit nooit in je afstudeerverslag!) Zelf ben ik dus inmiddels ruim een centimeter kleiner. Zou het – zo vroeg ik me tijdens het schrijven van dit hoofdstuk af – niet zo kunnen zijn dat de *meters* van tegenwoordig gewoon kleiner zijn? De euro is ook minder waard dan toen hij werd ingevoerd.

Taylor (1997) gebruikt in zijn boek het begrip ‘definitieprobleem’ (*problem of definition*). Hij geeft een voorbeeld van een timmerman (v/m) die wil meten hoe hoog een deur is. Dat doet hij of zij bijvoorbeeld met een rolmaat of duimstok, en die kun je vrij aardig aflezen. Op millimeters misschien wel. Maar het probleem met die deur is niet zozeer het aflezen en de eventuele onbetrouwbaarheid van datgene waarmee er wordt gemeten, maar de deur zelf. Die is nooit volmaakt recht en in oude huizen vaak zichtbaar scheef. Wat bedoel je dan precies met *de hoogte*? Het gemiddelde zou een aardig uitgangspunt kunnen zijn, maar de maximale hoogte is misschien wel nuttiger.

Zo’n definitieprobleem gaat nog een stap verder dan alleen scheef zijn. Er is niet alleen variatie in de *ruimte*, maar ook in de *tijd*. Zelfs als de deur niet meetbaar scheef is of als je besluit het midden van de deur aan te houden, kan bijvoorbeeld vocht ervoor zorgen dat de deur uitzet of krimpt.

Bij de meting van de lengte van een mens speelt dit soort verschijnselen ook een rol. Ons gewicht (natuurkundigen zouden zeggen: onze massa) varieert in de loop van een dag, van het jaar en van de jaren. Maar ook onze lengte is niet het hele etmaal door constant. Dat heeft te maken met de tussenwervelschijven die een mens in zijn of haar rug heeft. Terwijl de aarde om haar as draait, drukken die tussenwervelschijven wat in elkaar. Doordat een mensens lichaam er 23 heeft, kan dat over je hele rug genomen wel om een paar centimeter gaan. ’s Morgens ben je langer dan ’s avonds. Tenzij je dag- en nachtritme niet zo is dat je ’s nachts op bed ligt.

Wat voor die timmerman met z’n deur en die zeventienjarige jongen en zijn keuring voor militaire dienst geldt, is van toepassing op veel metingen. Ook als er geen definitieprobleem is, hebben we vrijwel altijd te maken met een *onzekerheid*. Over die onzekerheid gaat dit boek.

Onzekerheden noteren we met een plus-min symbool en een kleine delta. Mijn lengte is gelijk aan: $1,71 \pm 0,02 \text{ m}$: $l = 1,71 \text{ m}$ en $\delta l = 0,02 \text{ m}$.

Het schatten (meten) van die lengte l “voelt” als makkelijker dan het schatten van δl . Niet veel mensen zullen beweren dat ik 1,29 meter ben, behalve dan dommeriken die een meetlat van drie meter hebben en die per ongeluk ondersteboven houden. Toch scheelt dat nog geen 25% met de lengte die ik zelf opgeef. Die militair van toen overdreef met z’n millimeters, maar er zijn heus wel optimisten die denken dat je iemand met een marge van minder dan een hele centimeter kunt meten. Pessimisten zullen drie centimeter misschien wel te weinig vinden. Dat scheelt dus al 50%, beide kanten op.

We hebben het in dit eerste hoofdstuk over *onzekerheid*. In dit geval een absolute onzekerheid, die dezelfde dimensie heeft als datgene wat we meten, namelijk een lengte in meters. Een onzekerheid kun je ook opgeven als relatieve onzekerheid (in procenten). Dat is bij de meting van je eigen lengte niet gebruikelijk. Voor onzekerheid kom je ook wel de term *fout* tegen (of in het Engels: *error*), maar dat kan voor verwarring zorgen. In de volgende hoofdstukken gaan we het ook nog regelmatig hebben over onnauwkeurigheid en imprecisie. Die twee samen vormen je onzekerheid.

Oefeningen

1. “Zou het niet zo kunnen zijn dat de meters van tegenwoordig gewoon kleiner zijn?” Is deze (niet serieus bedoelde) uitspraak een theoretisch mogelijke verklaring voor het verhaal hierboven over het verschil van lengtes toen en nu?
2. Wanneer is de meter voor het laatst veranderd?
3. Uit welke bron heb je het antwoord op de vorige vraag?
4. In dit hoofdstuk wordt gesproken over een afwijking van *twee keer tweemaal de standaardafwijking*. Is dat niet dubbelop?
5. “Dat “tweeënzeventig zes!” duurde, als je de galm niet meerekende, ongeveer 1,414213562 seconde.” Welke twee aspecten uit dit hoofdstuk zijn verwerkt in deze (ook al niet serieus bedoelde) uitspraak?
6. “Waarom zou je het erbij zeggen, als het vanzelf spreekt?”

2. Snelheidscontroles

Op de website van het Openbaar Ministerie staat de volgende tekst te lezen in het menu Home > Onderwerpen > Verkeer > Handhaving in het verkeer > Snelheid en te hard rijden > Marges en meetcorrecties.

Snelheidscontroles zijn een veelbesproken onderwerp. Als de politie een voertuig geflitst heeft, wordt er met twee zaken rekening gehouden:

1. Meetcorrectie

Om met 100% zekerheid te kunnen garanderen dat iemand te hard heeft gereden, trekt de politie van de gemeten snelheid een paar kilometer af. De meetcorrecties zijn:

- *bij minder dan 100 km/u: 3 kilometer meetcorrectie*
- *bij meer dan 100 km/u: 3 procent meetcorrectie*
- *Het Nederlands Meetinstituut ikt alle flitsapparatuur.*

2. Ondergrens voor bekeuren

Er geldt een ondergrens van 4 km/u voordat wordt bekeurd.

Wat hier staat, heeft te maken met de inhoud van *Meten & Weten*. De tekst roept bovendien wat vragen op. Om te beginnen de eerste zin.

Snelheidscontroles zijn een veelbesproken onderwerp.

In feite staat hier heel weinig. Dat het ‘veelbesproken’ is, zal waar zijn. Maar wat was de inhoud van die gesprekken? Is er soms discussie over de vraag of de controles eerlijk of nodig zijn? Vindt men dat er te veel of te weinig wordt gecontroleerd? Hoeveel mensen denken dat eigenlijk en *waarom* dan?

In een afstudeerverslag heb je als student een probleem als je zoiets opschrijft. Er zijn twee dingen mis met die stelling.

1. ‘Veelbesproken’ is te weinig kwantitatief en daardoor niet te toetsen.
2. Er staat geen bron bij.

Het zou kunnen dat de opsteller van deze tekst de inhoud van wat er ‘veelbesproken’ is bewust opengelaten heeft. Sommigen zullen de controles te streng vinden en anderen te mild. De ene persoon gelooft in *meten* = *weten* en dat je niet moet zeuren maar bekeuren. Een ander vertrouwt de techniek die erachter zit niet. Je kunt zomaar een hele verjaardagsvisite vullen en bederven met een discussie erover.

Over de eerste zin van de tekst over de meetcorrectie valt ook nog wel wat te zeggen.

Om met 100% zekerheid te kunnen garanderen dat iemand te hard heeft gereden, trekt de politie van de gemeten snelheid een paar kilometer af.

‘Garanderen dat iemand te hard heeft gereden’ is een onhandige woordkeuze. Tenzij de politie graag wil dat iedereen de maximumsnelheid overschrijdt; dan zijn er inderdaad garanties nodig.

Kun je iets ook garanderen met minder dan 100% zekerheid? Is het dan nog wel garanderen? Die vraag laten we nu rusten.

Verderop in dit boek zullen we zien dat we het in dit vakgebied zelden over 100% zekerheid hebben, vanwege het simpele feit dat er voor 100% zekerheid extreem grote marges of middelen nodig zijn.

In de tekst wordt er gesproken over een ondergrens van 4 km/u... Bedoelen ze daarmee dat je nooit wordt bekeurd als je langzamer dan 4 km/u rijdt? Met die snelheid halen voetgangers je in. Er zal wel een ondergrens aan de *overschrijding van de maximumsnelheid* bedoeld worden. Staat dat er?

Wat het OM hierboven probeert te zeggen, is dat die onzekerheid wordt gecompenseerd door een correctie.

De meetcorrecties zijn:

- *bij minder dan 100 km/u: 3 kilometer meetcorrectie*
- *bij meer dan 100 km/u: 3 procent meetcorrectie*
- *Het Nederlands Meetinstituut ikt alle flitsapparatuur.*

Voor een goede technicus doet dit pijn aan de ogen... Een correctie van 3 km op een snelheid van 100 km/u. Dat betekent dat je een *afstand* probeert af te trekken van een *snelheid*. Dat is net zoiets als twee appels verwijderen uit een doos eieren. Het kan niet en het slaat nergens op. Laten we de tekst eerst corrigeren en het voortaan zelf nooit meer fout doen.

De meetcorrecties zijn:

- *bij minder dan 100 km/u: 3 km/u meetcorrectie*
- *bij meer dan 100 km/u: 3 procent meetcorrectie*

Wat een goede technicus zich ook meteen afvraagt: en als je nu eens exact 100 km/u rijdt? Strikt genomen zegt de tekst daar niets over. Toch kun je het wel met gezond verstand invullen. Elk van beide regels geeft voor 100 km/u namelijk dezelfde correctie.

Wat wel grappig is: als je met een snelheid van 2 km/u wordt geflitst, dan reed je volgens de gecorrigeerde meting met de helft van dat slakkengangetje de andere kant op... Je reedt dan namelijk -1 km/u.

Strikt genomen *corrigeer* je niet voor de meting, maar houd je rekening met een *marge* omdat er een *onzekerheid* in de meting zit.

Eigenlijk zou de vraag moeten zijn:

Weten we voldoende zeker dat de bestuurder te hard reed?

Kennelijk hebben de makers van de flitspalen voldoende vertrouwen om ‘ja’ te zeggen als hun meetapparatuur aangeeft dat iemand meer dan 103 km/u reed op een weg waar je 100 km/u mag rijden.

Is het mogelijk dat je daardoor bestuurders *niet* bekeurt terwijl ze *wel* te hard reden. Jazeker! En naarmate je grotere marges neemt, wordt die kans groter.

We hebben hier te maken met vier mogelijkheden.

	bekeuring	geen bekeuring
Reed te hard	terecht	?
Reed niet te hard	onrechtvaardig	gewenst

Van dit kwartet aan opties hebben we er aan drie een waardeoordeel gehangen door dat in te vullen in het schema.

- Het is wenselijk dat niemand te hard rijdt en dat niemand een bekeuring krijgt. Eigenlijk zou je die hele flitspaal niet willen hebben.
- Het idee van bekeuringen is dat je iedereen een prent wilt geven die te hard rijdt. Vandaar het woord ‘terecht’.
- Maar wat nu als je keurig 97,6 km/u reed en je kreeg toch die boete? Dat is oneerlijk en moet voorkomen worden.
- Iemand reed te hard en wordt niet bekeurd. ‘Die heeft geluk gehad’, zal iemand zeggen, maar is dat ook zo? Stel dat je telkens wegkomt zonder bekeuring en doorgaat met gevaarlijk rijgedrag...

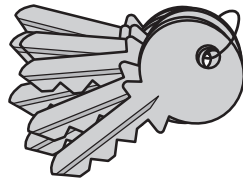
Oefeningen

1. Verbeter alle fouten in de cursieve tekst (punten 1 en 2) van de website van het Openbaar Ministerie.
2. Wordt iemand die 107,2 km per uur rijdt bekeurd als zij of hij geflitst wordt?

3. De sleutel

Toen ik in de vijfde klas van het vwo zat (1985/'86, Het Christelijk Lyceum Dordrecht) kregen we bij Wiskunde A een sommetje over statistiek.

Je hebt een sleutelbos met acht sleutels en eentje daarvan past op de deur waar je voor staat. Omdat je niet weet welke het is, ga je het domweg proberen. Hoe groot is de kans dat de vijfde sleutel die je probeert de juiste is?



Ik wist hoe je dat aanpakte. Je moest weten hoe groot de kans is dat een willekeurige andere sleutel paste. Dat was de kans dat de eerste sleutel niet goed was, en die moest je nou net hebben, want de som ging over het geval dat de *vijfde* goed was. Niet de eerste dus. Vervolgens berekende je van het setje dat overbleef opnieuw wat de kans op een verkeerde sleutel was, en dat vernieuwde je daar dan mee. Per stap werd het steeds waarschijnlijker dat je de juiste sleutel had, want alles wat niet paste, viel af. (Tenzij je dom of dronken was, maar we gingen uit van iemand die niet dezelfde sleutel nog een keer probeerde. “Kijken of-ie het inmiddels wel doet.”) Dat kon je netjes uitschrijven en dat deed ik ook. Het ging om de *vijfde* sleutel... We moeten dus eerst *vier keer* de kans nemen dat je een *verkeerde* sleutel te pakken hebt en dan *één keer* de kans op de *goede*. Nadat we vier keer misgegrepen hebben, blijven er nog vier sleutels over. Het eindigt dus met een kans van 1 op 4.

$$P = \frac{7}{8} \times \frac{6}{7} \times \frac{5}{6} \times \frac{4}{5} \times \frac{1}{4}$$

Eigenlijk wil je dit natuurlijk niet: een som uitrekenen waardoor je weet dat de *vijfde* sleutel in een bos van *acht* de juiste is. Straks komt er iemand die vraagt hoe het zit met de zesde sleutel in een bos van zeven. Of de zeventiende in een bos van honderdveertig. Of een speld in een hooiberg. Weet je wat, we gaan het in *Python* programmeren. Op een generieke manier, voor een willekeurig aantal sleutels.

Leuk! Een programmaatje maken in *Python*. Zullen we de variabelen *x* en *y* noemen? NEE! Die noemen we ‘aantal’ en ‘juiste’. Beter zou misschien wel zijn om ze ‘aantal_sleutels’ en ‘juiste_sleutel’ te noemen. Het belangrijkste is dat we volgend jaar onze eigen code nog snappen.

```

aantal = 8
juiste = 5
kans = 1

# bereken eerst de kans dat de eerste sleutels fout waren
for i in range(1, juiste):
    kans = kans * (aantal - i) / (aantal - i + 1)

# vermenigvuldig dan met de kans dat die gevraagde poging goed was
kans = kans * 1 / (juiste - 1)

# druk het resultaat af
print("De kans dat sleutel {} in een bos van {} de juiste is: {}".format(juiste, aantal, kans))

```

Terug naar het sommetje in 1986. Dat schreef ik met de hand uit, op een blaadje. En omdat ik dat blaadje niet meer heb, doe ik het nu maar opnieuw.

$$P = \frac{7}{8} \times \frac{6}{7} \times \frac{5}{6} \times \frac{4}{5} \times \frac{1}{4} = \frac{1}{8}$$

Bij het uitschrijven en vereenvoudigen van die breukberekening blijkt zo'n beetje alles tegen elkaar weg te delen! Alleen de noemer van die eerste factor (8) en de teller van die laatste factor (1) blijven over. Sterker nog: je zou dit voor willekeurige aantallen sleutels (N dus) kunnen doen.

De kans dat de vijfde sleutel de goede is, blijkt gelijk te zijn aan $1/N$. De kans voor een willekeurige andere sleutel is hetzelfde.

Ik herinner me nog de glimlach op mijn eigen gezicht destijds. (Of ik heb die er later bij verzonnen, dat kan ook.) Want wat had ik moeilijk gedaan over zo'n simpele vraag! De kans dat de vijfde sleutel aan een bos van acht de juiste was, is natuurlijk gewoon *een achtste*. Dat zie je in één keer als je je realiseert dat alle sleutels even waarschijnlijk zijn.

Om een of andere reden is dit sommetje in mijn geheugen blijven hangen. Net als die keer dat ik geen 10 kreeg voor mijn foutloos gespelde dictee. En dat ene moment waarop Hans per ongeluk een liniaal brak tijdens Teken en. Sommige dingen vergeet je nooit. Andere meteen nadat je ze gehoord hebt. Dan vergeet je het ook maar één keer trouwens.

De omslachtige oplossing staat tegenover de aanpak van het gezonde verstand. In dit boek geven we aan beide zaken aandacht. *Met en & Weten* gaat over onzekerheden en fouten. Daarbij wil je enerzijds dingen correct en formeel doen, en ben je anderzijds voortdurend bezig met de vraag of het niet anders en makkelijker kan.

De kans dat die gevraagde sleutel de juiste was, hebben we nu bepaald

- via een formele som,
- met gezond verstand,
- in een Pythonprogramma.

Er kwam telkens hetzelfde uit... De vraag is dan nog steeds of het antwoord juist is. Als er verschillende antwoorden uit waren gekomen, was er iets fout. Nu we drie keer hetzelfde hebben, *zou* het correct *kunnen* zijn.

De schrijver van dit boek is benieuwd naar wat de lezer dacht na het lezen van bovenstaand vraagstuk. Sommigen zagen in één keer dat het “gewoon” $1/8$ was. Anderen maakten diezelfde stomme som als die schooljongen in de jaren tachtig. Weer anderen lazen met glazige ogen verder. (Of niet.) Er zullen ook lezers zijn die zelfs nu nog geen idee hebben dat en waarom het inderdaad een kans van 1 op 8 is.

Alle sleutels zijn even waarschijnlijk de juiste. Zo simpel is het.

Misschien zegt iemand: de kans is nul. Volgens de Wet van Murphy is het altijd pas de laatste sleutel die past.

Oefeningen

1. Hoe groot is de kans dat bij de sleutelbos uit dit hoofdstuk één van de eerste vijf sleutels die je probeert de juiste is?
2. Welke redeneerfout maakt iemand die denkt dat de kans dat die vijfde sleutel de juiste blijkt 50% is?
3. Bij een hbo-opleiding met een instroom van zo'n 120 studenten per jaar en een uitvalspercentage van bijna 60% staan er 389 studenten ingeschreven. Hoe groot is de kans dat twee of meer studenten van deze opleiding op dezelfde dag jarig zijn?

Een man loopt 's nachts op straat te zoeken naar zijn huissleutel, die hij heeft laten vallen. Hij tuurt en tiert en loert en vloekt en baalt en haalt mismoeedig zijn schouders op.

“Die ga ik nooit meer vinden in deze vieze, vuile, verrekte...”

Nog net voordat hij een héél lelijk woord wil gebruiken, hoort hij achter zich de stem van een goede vriend.

“Wat is er met jou aan de hand, joh?”

“Ik heb mijn sleutel laten vallen en kan hem niet vinden.”

“Zal ik je helpen zoeken?”

“Nou, als je dat zou willen doen!”

“Waar heb je 'm ongeveer laten vallen?”

“Een meter of zes die kant op...”

“Maar waarom zoek je dan niet daar?”

“Het is verderop nogal donker.

Hier heb ik meer licht van die lantaarnpaal.”



Dit verhaal is, zoals je waarschijnlijk al vermoedde, niet echt gebeurd. Nou ja, misschien is het wel echt gebeurd, maar mij is het niet bekend.

Het is een van mijn favoriete moppen. Niet zozeer omdat hij leuk is, maar vooral omdat er iets van de werkelijkheid in zit. Als mensen iets niet kunnen vinden of niet weten, zijn ze geneigd om hun toevlucht te nemen tot iets waar ze wél mee uit de voeten kunnen. We zoeken liever in het licht (figuurlijk gesproken) naar iets wat er helemaal niet is, dan in het donker (ook figuurlijk) op een terrein waar we te weinig verstand van hebben.

Voor meten en weten geldt dat er heel veel is wat we *niet* weten of niet kunnen meten. De verleiding is groot om dan maar dingen te gaan aan-nemen, buiten beschouwing te laten of domweg te negeren.

4. Doellijntechnologie

Feyenoord heeft in Rotterdam een enerverende eredivisietopper tegen PSV met 2-1 gewonnen. Er zal nog lang nagepraat worden over het winnende doelpunt, dat acht minuten voor tijd op basis van doellijntechnologie werd toegekend. Nadat Jan-Arie van der Heijden na veel duw- en trekwerk los was gekomen van Héctor Moreno, zag hij zijn kopbal op de lijn gestopt worden door PSV-doelman Jeroen Zoet. Feyenoord claimde tevergeefs een goal: het horloge van scheidsrechter Bas Nijhuis (waarop doellijntechnologie zichtbaar is) gaf geen doelpunt aan. Bij het opstaan rolde Zoet de bal echter naar achter en daarbij passeerde die alsnog de doellijn. Toen pas ook ging het horloge van Nijhuis af. Het ging echter wel om millimeters.

Bron: NOS, 26 februari 2017

Ging het inderdaad om millimeters? Het artikel zegt van wel, maar het horloge van Bas Nijhuis deed daar geen uitspraak over. Daar stond het Engelse woord GOAL in hoofdletters, omdat het systeem een doelpunt vaststelde.

Als je goed over dingen nadenkt, kloppen ze zelden. (Zelfs op deze zin is al het nodige af te dingen.) Dat horloge kan natuurlijk wel GOAL roepen, maar dát is niet wat er vastgesteld wordt. Een overtreding of buitenspelsituatie kan er namelijk voor zorgen dat het wijze waarop de bal de doellijn passeerde niet rechtmatig was. De doellijntechnologie ziet alleen de bal en de lijn. Niet het spel.

Wat zou er gebeurd zijn als dat doelpunt in Rotterdam-Zuid niet opgemerkt en toegekend was? Dan was het (op dat moment) geen 2-1 geworden en wellicht zelfs 1-1 gebleven. Geen winst voor de thuisclub (en op dat moment koploper), maar een gelijkspel.

De eindstand in de competitie van het seizoen 2016/2017, waarin deze wedstrijd plaatsvond, zag er als volgt uit.

			W	V	G	+	-	+/-	pt.
1	Feyenoord	34	26	4	4	86	25	61	82
2	Ajax	34	25	3	6	79	23	56	81
3	PSV	34	22	2	10	68	23	45	77

Alle clubs hadden 34 wedstrijden gespeeld. Iedere overwinning (kolom W) leverde drie punten op en elk gelijkspel (kolom G) was één punt waard. Het aantal doelpunten voor en tegen en het doelsaldo (verschil tussen voor en tegen) was ook berekend. Hoe zou dit lijstje eruit hebben gezien als het 1-1 was gebleven? De nummer 1 had dan één voordoelpunt en één overwinning minder gehad en de nummer 3 kreeg juist één tegendoelpunt minder en leed één nederlaag minder. Beide clubs hadden dan één keer meer gelijkgespeeld.

De eindstand had er dan als volgt uitgezien.

			W	V	G	+	-	+/-	pt.
1	Ajax	34	25	3	6	79	23	56	81
2	Feyenoord	34	25	4	5	85	25	60	80
3	PSV	34	22	1	11	68	22	46	78

Ach, het is maar een spelletje. Toch zou er in Amsterdam en Rotterdam anders naar dit lijstje gekeken zijn...

Ondertussen is het maar de vraag of het 1-1 gebleven zou zijn als dat doelpunt niet gevallen was. In de laatste minuten kan het soms raar lopen als de belangen groot zijn. En ook die tien wedstrijden die er nog te spelen waren, zouden wellicht anders verlopen zijn als deze uitslag anders was geweest.

‘Wat er gebeurd zou zijn als’ is een uitspraak die zelden te bewijzen of te controleren is. Had je geen leuke band gehad, dan was je toch nog op tijd gekomen voor het tentamen. Maar misschien was je dan wel geschept geweest door een motorfiets die te hard reed en uit de bocht vloog. Meten is weten en zodra je iets meet waarvan je weet dat het van belang is, doet dat iets met je. Je gaat anders voetballen, wat harder fietsen, meer of minder rekening met dingen en mensen houden.

Wat was eigenlijk de onzekerheid in het meetsysteem? De voetbalbond KNVB claimde destijds dat de foutmarge van doellijntechnologie ongeveer een millimeter is. Peter de With, hoogleraar videotechnologie van de TU Eindhoven, trok dat in twijfel. Hij stelde dat de foutmarge eerder een centimeter was. ‘In een laboratoriumomgeving kun je het testen, maar dan zijn de doelpalen recht en is de doellijn nauwkeurig’, zei hij tegen Omroep Brabant. ‘In de werkelijkheid is dat anders. Dan heb je met grotere marges te maken. Een doellijn wordt getrokken met een kalkmachine, op

gras dat niet glad is en waarvan de sprietjes alle kanten op staan. Doelpalen zijn ook nooit honderd procent verticaal en recht. Daardoor is de berekening van de doelpositie veranderlijk, waardoor de doellijn in het model verschuift. Dit geeft foutmarges in de exacte positie van de doellijn.'

Een definitieprobleem dus. Waar staat het doel precies en hoe loopt die lijn? Niet zozeer het *meetapparaat*, maar *datgene wat je meet* is onzeker. Daarop komen we terug in hoofdstuk 15.

Een interessante vraag is wat je als technicus doet met die onzekerheid. In het geval van doellijntechnologie: als de bal een halve centimeter over de doellijn 'gemeten' wordt, en je onzekerheid is een hele centimeter, telt het doelpunt dan wel of niet? Het kampioenschap kan er zomaar van afhangen...

Oefening

Je werkt als mechatronicus voor de KNVB en krijgt de opdracht om een nieuw systeem voor doellijntechnologie te realiseren. Het systeem dat je ontwerpt, kent een onzekerheid van 2 cm. Je zou daarvoor kunnen corrigeren door de software zo te schrijven dat het horloge pas 'GOAL' roept als de bal minstens 2 cm over de lijn is. Je kunt ook *niet* corrigeren en de onzekerheid voor lief nemen. Of nog iets anders doen.

Maak een keuze tussen de mogelijke oplossingen en motiveer die.

5. Onzeker zijn



Het “plaatje” hierboven heeft geen figuurnummer, omdat het geen figuur is. Het is een *tekst*. Het bestaat uit alleen maar Unicode tekens. Dat kun je controleren door je cursor achter elk van de tekens te zetten in *Word*, en dan tekens op Alt+X te drukken. Voor het eerste karakter in “_(\ツ)_” vind je dan 00AF en voor het tweede 005C. Er zit er eentje bij die wat raar is in de Nederlandse taal: het teken ツ is “gewoon” Unicode 30C4, het Japanse teken Tsu. Volgens *Wikipedia* dankt dat teken in het Westen zijn bekendheid hieraan.

The shrug gesture is a Unicode emoji included as U+1F937. The shrug emoticon, better known as the shruggie, made from Unicode characters, is also typed as “_(\ツ)_”, where “ツ” is the character tsu from Japanese katakana.

Het schijnt trouwens dat de Chinese taal duizenden karakters heeft die eigenlijk allemaal plaatjes en tekeningen op zich zijn. Je zou dus verwachten dat onzekerheid in het Chinees gewoon één symbool is. Toch zegt Google Translate dat het er vier zijn: 不安全感.

En dan nog een vraag: is schouderophalen (Engels: *shrug*) wel hetzelfde als onzeker zijn? Zeggen dat je iets niet weet of dat het je niet uitmaakt, is geen uiting van onzekerheid. Het kan ook onverschilligheid aanduiden.

Meten & Weten gaat vaak over techniek & wiskunde, metingen & sommen. Soms is het handig om even na te denken over al het *menselijke* daarachter. In het vorige hoofdstuk lasen we over meetfouten: “Iets wat de hele tijd gelijk blijft, komt niet over als onzeker.” Dat lijkt een beetje op hoe sommige mensen in elkaar zitten: nooit een fout erkennen, altijd vol blijven houden dat je gelijk hebt. (En als je dan tóch een keer moet toegeven dat het niet A is maar B, vanaf dat moment doen alsof je altijd al B gevonden hebt...)

Doen alsof je onzekerheid kleiner is dan ze is, straalt echter geen zelfvertrouwen uit maar *angst*. Het is ook niet slim, maar *dom*. Ongelijk hebben is namelijk alleen maar *winst*: je aantal onjuiste meningen neemt met eentje af en je aantal juiste meningen neemt met evenveel toe. Het verdedigen van iets wat fout is, is ook nog eens veel lastiger dan concluderen dat je ongelijk hebt. Anderen vinden je meestal ook sympathieker als je ze gelijk geeft. (Daarin moet je ook weer niet doorschieten, want gelijk geven als ze het niet hebben, heeft ook weer nadelen.)

Voor metingen en meningen geldt niet alleen dat maar één letter verschillen. Sommige mensen veranderen nooit van inzicht of mening. “Iets wat de hele tijd gelijk blijft, komt niet over als onzeker.” Mensen die sterk vasthouden aan hun mening, komen ook niet over als onzeker. Soms zijn ze dat juist wél. Ze willen de fout niet toegeven, omdat iedereen dan ziet dat ze een “misser” begaan hebben. Maar iemand die een fout maakt én blijft volhouden dat het goed is, komt uiteindelijk dommer over dan iemand die het inziet en dat laat merken. Glimlachend, want fouten maken hoort gewoon bij mens-zijn.

Oefeningen

Deze oefeningen moet je vooral niet doen, want ze hebben geen enkel nut. Alleen de eerste is zinnig. Sla de rest dus over, lees ze niet en ga naar het volgende hoofdstuk. Zonde van je tijd om 2 t/m 4 te doen.

1. Als jij een foute gedachte over iets hebt, ben jij dan fout of is alleen die gedachte fout?
2. Doe een eenvoudige dyslexie-test door te kijken of je het verschil ziet tussen de woorden *meting* en *meting*.
3. Maak (met Inkscape of een ander programma) een vectorgebaseerd teken dat “schouderophalen” betekent.
4. Vraag je eens af waarom je deze oefeningen gelezen (en misschien wel gedaan) hebt, terwijl er boven de opdracht stond dat ze geen enkel nut hebben.

6. W'rom?

Als een docent het over de WHW heeft, bedoelt hij of zij meestal de *Wet op het hoger onderwijs en wetenschappelijk onderzoek*. Zelf denk ik ook vaak aan een oud-collega die Deze Drie Letters groot opschreef tijdens het bespreken van het conceptverslag van een afstudeerder.

“In een goed verslag staat iets over het Wat, het Hoe en het Waarom.

Wat een student gedaan heeft, wordt eigenlijk altijd wel gezegd.

Hoe hij of zij dat gedaan heeft meestal ook wel.

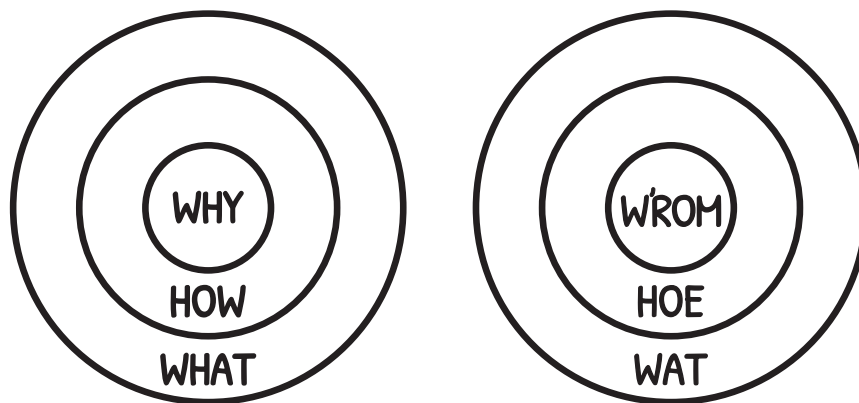
Waarom het zó gegaan is en niet anders, ontbreekt vaak.

Terwijl dat nu juist is waar je als examinerator in geïnteresseerd bent.

Je wilt kunnen beoordelen of de student in staat is om de juiste keuzes te maken en goede argumenten te gebruiken.”

Waarom? In *Start with Why: How Great Leaders Inspire Everyone to Take Action* van Simon Sinek (2009) gaat het ook om die drie woorden. Daar heeft het geen betrekking op afstudeerverslagen, maar op het motiveren van medewerkers door leidinggevenden. Als iemand te horen krijgt *wat* hij moet doen, is dat gewoon een opdracht. Misschien hoort hij ook nog *hoe*. Als hij weet *waarom* het nodig is (en doel en nut inziet), doet hij het meestal beter.

Sinek gebruikt in zijn boek drie concentrische cirkels.



Figuur 6.1 The Golden Circle van Sinek, en de w'rom-variant erop

Altijd doen! Een plaatje maken en dan iets met rondjes, of blokken met pijlen – bij voorkeur in een driehoek. En kleurtjes, als dat kan. Dan lijkt het allemaal heel wijs en belangrijk.

Het Engels heeft als voordeel dat *What – How – Why* (of andersom) lekker strak en kort en éénlettergrepig is. Hollanders hebben met ‘waarom?’ net als de Duitsers met ‘weshalb?’ voor dat derde woord zes of zelfs zeven letters voor nodig. Fransen moeten het doen met ‘pourquoi?’, wat naar twee

woorden neigt. In het Spaans *zijn* het er zelfs twee: ‘¿por qué?’ In Italië zeggen ze ‘perché?’ Zowel in het Deens als het Noors: ‘hvorfor?’ In het Zweeds: ‘varför?’ In *Metén & Weten*: ‘w’rom’?

Waarom? Om wat? Voor wat? Met welk doel?

De voornaamste doelgroep van *Metén & Weten* wordt gevormd door studenten die over vijf maanden gaan afstuderen. Of over een jaar, of anderhalf. Zulke studenten doen zichzelf en de lezer van hun verslag een plezier als de antwoorden op de vraag ‘waarom’ er duimendik bovenop liggen.

W’rom staat er eigenlijk iets over WHW in *M&W*? Leuk dat je het vraagt! Het werkt nu al, dat hameren op ‘w’rom’. Hier is het antwoord.

Bij metingen heb je ook altijd met een *reden* of *doel* en *aanleiding* te maken. Wat je gemeten hebt, wordt uitgedrukt in een grootte met een onzekerheid.

Hoe je dat gedaan hebt, wordt soms ten onrechte niet vermeld, maar is nou juist waar die onzekerheid vandaan komt.

Waarom (of: w’rom) je het wou weten is – als het goed is – doorslaggevend geweest voor de manier van meten die je koos. Het w’rom bepaalt het hoe.

Een tijd die je heel nauwkeurig en precies wilt kennen, kun je meten met een atoomklok. Een stopwatch kan ook, als het minder krap komt. Hardop of in gedachten tellen in een seconde-achtig ritme, is soms ook toereikend.

Maar w’rom zou je een atoomklok kopen als een dom klokkie ook genoeg is? En w’rom zou je tellen als er een stopwatch op de telefoon in je zak zit? Of een horloge om je pols, maar wie heeft dat tegenwoordig nog?

Vroeger had je elpees. De eerste die ik ooit kreeg, was *10 jaar André van Duin*. Voor m’n zevende verjaardag, zomer 1976. Daarop stond onder andere *Het Bananenlied*, dat draaide om de vraag: waarom zijn de bananen krom? Hieronder het begingedeelte van de liedtekst, die al heel snel het antwoord op die vraag verradt.

Waarom? Wa-ha-harom?

Waarom zijn de bananen kro-hom?

W’rom zijn de b’na zijn de b’na zijn de b’na zijn de b’na zijn de b’na-ha-nen krom?

Als je ze rechttop zet dan vallen ze om.

7. Model & Werkelijkheid

“All models are wrong, but some are useful.”

George Box

Bovenstaande tegeltjeswijsheid is afkomstig van George E. P. Box en de eerste keer dat die zwart op wit kwam te staan was in 1976, in een wetenschappelijke publicatie in *the Journal of the American Statistical Association*.

Dat woord *wrong* doet denken aan fouten die je hebt bij metingen. En daar zit wel een overeenkomst in. Bij metingen zeggen we dat we te maken hebben met een *fout*, hoewel we het in dit boek meestal (en liever) over een *onzekerheid* hebben. Dat is niet hetzelfde.

De meting die je doet, komt niet exact overeen met de werkelijkheid. Ook als je niets verkeerd doet. Een telling kan exact kloppen, maar de meting van een lengte, massa, tijd, spanning, druk of wat dan ook, heeft altijd een beperking in zich. Je zou de uitspraak van Box een beetje kunnen aanpassen.

“All measurements are wrong, but some are useful.”

Het *some* is in beide uitspraken een beetje cynisch of komisch bedoeld. Het is niet zo dat er maar een paar nuttige modellen of metingen bestaan. Maar het is belangrijk om te zien dat ze niet een exacte weergave van de werkelijkheid zijn.

Een leraar van me op de middelbare school zei het anders.

“Alle modellen zijn minder dan de werkelijkheid, behalve fotomodellen. Die zijn méér dan de werkelijkheid.”

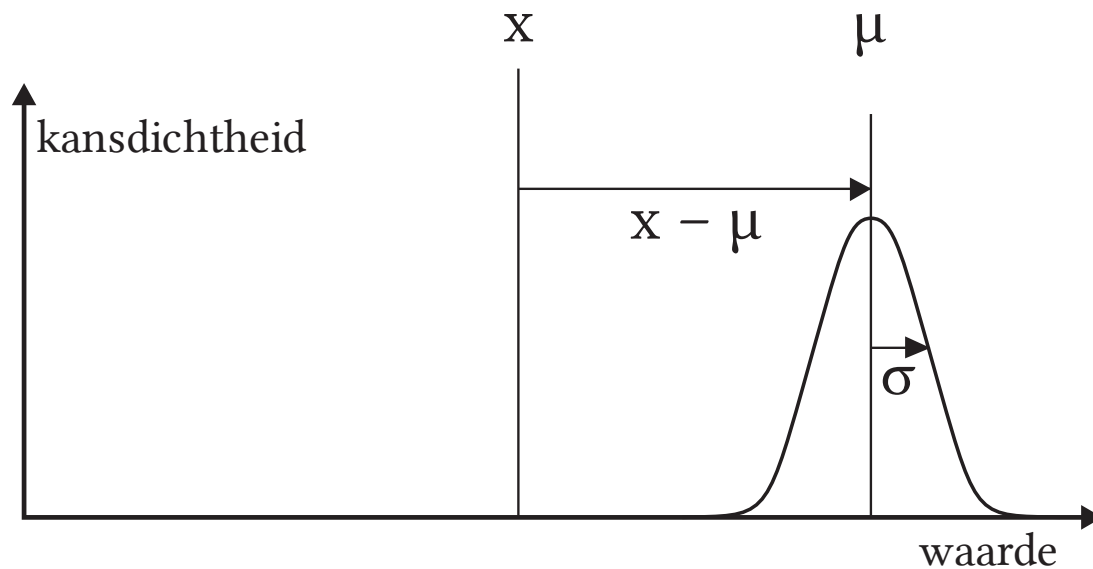
Soms tref je in *Metten & Weten* hoofdstuktitels aan met twee woorden en een & ertussen. Beide helften beginnen dan met een hoofdletter. In dat geval worden er twee begrippen naast elkaar gezet die écht iets anders betekenen en goed uit elkaar moeten worden gehouden.

Het Ene & Het Andere is niet hetzelfde, maar enigszins anders. Of heel anders. Dat zit 'm in de verschillen...

Fouten & Onzekerheden

8. Model van een meting

Een klassiek plaatje bij het uitleggen van hoe systematische en toevallige fouten in elkaar zitten, is weergegeven in Figuur 8.1.



Figuur 8.1 Meting met systematische fout en normaal verdeelde toevallige fout

Deze figuur is overgenomen van *Wikipedia*, met aanpassingen: de woorden *juistheid* en *precisie* zijn vervangen door symbolen.

De waarde die je wilt meten (de werkelijke waarde of de *true value*) is x .

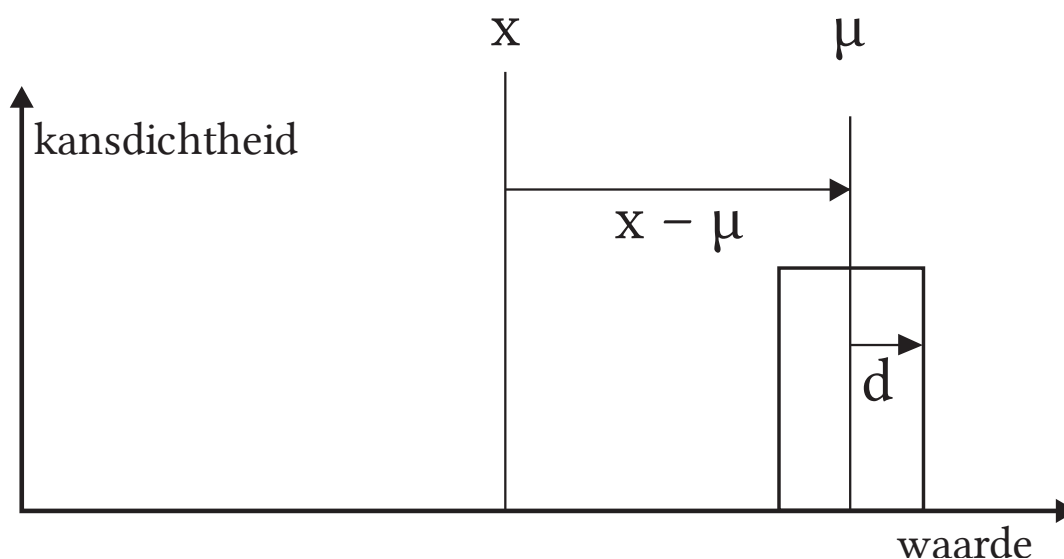
De meting zelf is een *kansvariabele*, normaal verdeeld met een gemiddelde μ en een standaardafwijking σ . Als je oneindig vaak meet en vervolgens de resultaten middelt, wordt het gemiddelde van je metingen gelijk aan μ . Als je N keer meet, is je meetresultaat (het gemiddelde van je metingen) een normaal verdeelde kansvariabele met een gemiddelde μ . De standaardafwijking van dat gemiddelde is

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{N}}$$

Je ziet in deze formule dat voor zeer grote aantallen metingen de spreiding van het gemiddelde heel klein wordt. Het gemiddelde van je steekproef komt steeds dichterbij het gemiddelde van de verdeling te liggen.

We gaan er in dit soort figuren al snel van uit dat een meting beschouwd kan worden als een normaal verdeeld proces met een gemiddelde μ en een standaardafwijking σ . Dat is een nogal oppervlakkige aanname. Het zou niet

vreemd zijn als die verdeling uniform was, bijvoorbeeld bij een digitaal meetapparaat dat systematisch een te hoge waarde aangeeft én een toevallige fout heeft die afhangt van de decimalen die je niet kunt aflezen.



Figuur 8.2 Meting met systematische fout en uniform verdeelde toevallige fout

Laten we voor het gemak eens uitgaan van een meting aan een dimensieloze (vage, abstracte) grootheid x met eenheid *nono* (symbool: n), waarvoor geldt

$$x = 1,98 n$$

Het meetresultaat (gemiddelde van tien metingen) dat we vinden is

$$x_{m1} = 2,01 \pm 0,04 n$$

De absolute onzekerheid in dit meetresultaat is 4 cn (centinono).

De relatieve onzekerheid in dit meetresultaat is 2%.

Weten we nu hoe groot het verschil is tussen x en μ ? Nee. In de tekst hierboven is *gegeven* dat $x = 1,98 n$. In een ‘echte’ situatie weet je dat meestal niet. Beter gezegd: weet je dat vrijwel nooit. Je kent soms een *geaccepteerde waarde*, maar slechts zelden de *werkelijke waarde*. Waarom zou je nog meten als je de waarde al kent?

Kennen we μ eigenlijk wel? Nee. We hebben een waarde x_{m1} gevonden. Die x_{m1} is het steekproefgemiddelde $\bar{x}$ van deze ene steekproef van tien metingen. Doe je nog een keer tien metingen, dan vind je een x_{m2} die – als de toevallige fout niet nul is – anders zal zijn. Naarmate de standaardafwijking σ groter is, zal de spreiding in x_{m1} , x_{m2} , x_{m3} , etc. groter zijn.

Kennen we σ dan wél? Nee. Die is ook al niet bekend in praktische situaties. We kunnen wel een schatting ervan maken door vaak te meten en dan de standaardafwijking van de steekproef te bepalen. Als je oneindig vaak meet, wordt de standaardafwijking die je vindt gelijk aan de ‘echte’ σ .

Is er eigenlijk wel een ‘echte’ σ ? Nee. We veronderstellen dat er een normale verdeling ten grondslag ligt aan die metingen, maar dat is slechts een model.

“All models are wrong, but some are useful.”

George Box zei dit, en herhaalt het nog maar eens.

Je raadt het al: er is ook geen ‘echte’ μ . Maar het gemiddelde van de veronderstelde (verzonnen, bedachte, er met de haren bijgesleepte) verdeling kun je wel aardig schatten door het gemiddelde van veel steekproeven te nemen.

Laten we niet al te principieel zijn en toch maar doen alsof er wél een echte μ en σ bestaan. Dan kennen we die nog lang niet. En de echte x kennen we alleen maar doordat we die zelf verzonnen hebben in dit hoofdstuk. Op grond van de metingen weten we alleen dit:

$$x_{m1} = 2,01 \pm 0,04 n$$

Hoe is die onzekerheid van $0,04 n$ eigenlijk bepaald? Daar hebben we het nog niet over gehad. Er is een constant deel en een toevallig deel in die onzekerheid. Het constante deel kan niet uit de metingen worden afgeleid, maar zou gegeven kunnen zijn in specificaties van de meetapparatuur. Het toevallige deel kan volgen uit die specificaties. Het toevallige deel kan óók worden gevonden door de metingen te herhalen en te kijken wat de spreiding is. Door oneindig vaak te herhalen, kun je de kansverdeling bepalen.

Er is een belangrijk aspect in het soort plaatjes aan het begin van dit hoofdstuk dat meestal onbenoemd blijft. Die plaatjes geven namelijk de verdelingsfunctie van één bepaalde meting voor één bepaalde grootheid.

We hadden gevonden als meetresultaat:

$$x_{m1} = 2,01 \pm 0,04 n$$

De absolute onzekerheid was 4 cn en de relatieve onzekerheid was 2%. Hoe groot zullen nu de absolute en relatieve onzekerheid zijn in een grootheid die twee keer zo groot is? En wéér komt dat inmiddels irritante antwoord: dat weten we niet.

Terug naar Figuur 8.1 waar we mee begonnen. Zou het waar zijn dat onze meting echt normaal verdeeld is? Nee. Als je een tijdsduur meet en je komt uit op gemiddeld 1,5 s en een standaardafwijking van 0,3 s, dan zou er bij een normale verdeling een (kleine) kans zijn dat die tijd negatief is. En negatieve tijdsduren bestaan niet.

Dat zou mooi zijn, zeg! Een negatieve tijdsduur... Je zit op de fiets naar school en komt waarschijnlijk vijf minuten te laat voor het tentamen. Je snijdt een flink stuk af via een route waar je *een negatief kwartier* over doet... Dan heb je toch nog tien minuten speling!

Die verdeling is dus niet ‘echt’ normaal. En als die verdeling dat wel was, kenden we nog steeds die μ en σ niet.

Er is nog meer aan de hand. Die onzekerheid kon per gemeten waarde wel eens verschillen. Op een tijdsduur van 1,5 s zit een onzekerheid van 0,3 s. Dat is 20%. Zit er dan op een tijdsduur van 2,5 s óók een onzekerheid van 0,3 s? Of óók een onzekerheid van 20% en dus 0,5 s? Misschien is het ene waar, misschien het andere... en misschien ook allebei niet.

Het zou kunnen dat de onzekerheid op die meting gelijk is aan 0,15 s + 10%. Het zou zelfs zo kunnen zijn dat er een kwadratisch of logaritmisch verband is tussen het gemiddelde en de standaardafwijking. Theoretisch kan er heel veel. In praktijk zijn we al blij als het ongeveer normaal (of uniform) verdeeld is en we daar wat getallen bij kunnen schatten.

Wat is hier mis met het model dat we gebruiken? Op zich niet zo veel, maar je moet je altijd realiseren dat die normale verdelingen met z'n μ en σ *slechts een model zijn*. En dat George Box gelijk had dat elk model fout is.

Oefeningen

1. Als George Box gelijk heeft dat elk model fout is, waarom gebruiken we dan modellen?
2. Als je een tijdsduur meet en je komt uit op gemiddeld 1,5 s en een standaardafwijking van 0,3 s, dan zou er bij een normale verdeling een (kleine) kans zijn dat die tijd negatief is.
Hoe groot is die (kleine) kans?

Voor opgave 2 heb je de kennis nodig die, als je die nog niet hebt, pas wordt behandeld in de hoofdstukken 87 en 88.

9. Beste schatting ± onzekerheid [eenheid]

Als je de eenheid weglaat, is je meetresultaat FOUT.

Als je de onzekerheid weglaat, weet je NIETS.

Als iemand zegt: “Ik ben Twee Zeven”, wat weet je dan?

Niet veel, maar je kent wel de voor- en achternaam van de heer of mevrouw T. Zeven. Er waren volgens de [Nederlandse Familienamenbank](#) in het jaar 2007 in Nederland 72 mensen met de achternaam “Zeven”. In de hoofdstad waren er precies zeven Zevens.

Misschien bedoelde degene die zei “Ik ben Twee Zeven” wel iets anders, namelijk “Ik ben twee meter en zeven centimeter lang.”

In het dagelijkse leven doen we dat vaak zo. “Het wordt morgen dertig graden” kan betekenen dat de *verwachting* volgens *het KNMI* is dat het morgen dertig graden *Celsius* wordt. In wetenschap en techniek (en op de hogeschool en op het tentamen en in je verslag en je latere baan) leggen we de lat hoger. Eenheden mag je niet weglaten. Onzekerheden ook niet.

Natuurlijk weet je in praktijk wel iets van iemand die zegt “twee zeven” te zijn. Dat zijn heus wel meters. En ook als je de onzekerheid niet kent: die persoon kan best wel “twee zes” of “twee vijf” of “twee vier” zijn, maar heus geen “één vijfentachtig”. Ook al is de onzekerheid niet vermeld: daar heb je bij een lengtemeting wel een gevoel bij. De onzekerheid is zeker groter dan een millimeter en zeker kleiner dan een decimeter.

Toch is de norm in dit boek strenger.

Als je de eenheid weglaat, is je meetresultaat FOUT.

Soms kan het echte verwarring geven, soms is het alleen maar onhandig en snapt de ander je ook wel. Maar het moet gewoon echt kloppen wat er staat en dat doet het niet zonder die eenheden.

Dat de ander je ook wel snapt, is overigens niet altijd waar. De [Mars Climate Orbiter](#) ging op 23 september 1999 verloren doordat de berekeningen voor de stuurmotoren werden uitgevoerd in pond-seconden in plaats van de newton-seconden die de software aanhield. De missie van 300 miljoen dollar ging de mist in. Maar het leverde wel een mooie illustratie op om uit te leggen dat correcte eenheden soms best belangrijk zijn.

Als je de onzekerheid weglaat, weet je NIETS.

Stel dat de spanning op de pin van een microcontroller doorslaggevend is voor de vraag of het brandalarm moet afgaan. De spanning kan 0 V zijn of 5 V. Maar in praktijk natuurlijk nooit helemaal exact, en dus wordt alles onder de 2,5 V beschouwd als 0 V en wordt alles boven de 2,5 V beschouwd als 5 V.

Je doet een meting en je meet: “Twee zeven”.

Dat is fout, want je moet een eenheid opgeven: 2,7 V.

Maar hoe zit dat met de onzekerheid? Als die 0,3 V is, zou de werkelijke spanning ook 2,4 V kunnen zijn. En als die 1,4 V is, misschien zelfs maar 1,3 V... Dat is heus wel minder dan 2,5 V. En om het nog erger te maken: misschien is de onzekerheid wel tien keer zo groot als die 0,3 V. Dan zou de spanning op de pin zelfs negatief kunnen zijn.

We meten dus “twee zeven” en interpreteren het als een hoge spanning. Misschien was het wel negatief.

Terug naar die lengtemeting uit hoofdstuk 0. Hoe noteer je het correct? Laten we aannemen dan de onzekerheid anderhalve centimeter is. Dan geef je de *beste schatting* op met daarachter een *plusminusteken* en de onzekerheid, gevolgd door de eenheid.

$$l = (1,726 \pm 0,015) \text{ m}$$

De haakjes mag je weglaten, omdat het vanzelf spreekt dat de eenheid betrekking heeft op beide getallen.

Volgens de [GUM](#), de *Guide to the expression of uncertainty in measurement*, is de officiële notatie anders.

$$l = 1,726 (15) \text{ m}$$

Het getal tussen ronde haken geeft de onzekerheid aan en moet worden gelezen als de onzekerheid in de laatste cijfers. In dit boek volgen we die Meestal niet, omdat die niet gebruikelijk is, al is het dan de officiële standaard.

Oefeningen

Als de hierboven genoemde gemeten spanning op het display werd afgelezen als 2,697 V en de onzekerheid was 0,3 V...

1. Hoe noteer je dan de spanning volgens de conventie in dit boek?
2. Hoe noteer je dan de spanning volgens de richtlijnen van de GUM?

10. Systematisch & Toevallig

Als we ervan uitgaan dat een grootheid goed gedefinieerd is en dat we niets verkeerd doen en dat er ook geen rare (of grote) spreidingen zitten in datgene wat we meten, dan hou je eigenlijk maar twee componenten over in je onzekerheid.

- de onzekerheid als gevolg van een toevallige fout
- de onzekerheid als gevolg van een systematische fout

Het toevallige deel verschilt per meting. Het systematische deel is altijd hetzelfde: je meet altijd te veel of te weinig en de hoeveelheid varieert niet.

Een misverstand dat je soms tegenkomt, is dat alleen het toevallige deel als ‘onzekerheid’ wordt gezien.

Laten we om dit te illustreren nog eens naar zo’n beginzin kijken, uit het boek van Barford (1985).

Words like error, accurate, truth, probable, likelihood, are part of everyday language. For most of us their meanings may not be mathematically precise, but the ideas behind them reflect experience we have gained from an early age.

Barford kiest zijn uitgangspunt bij de alledaagse ervaring van begrippen. Op zich niets mis mee. Maar dan gaat hij verder.

One way of approaching the subject matter of this book is to step aside from this experience and to talk of tossing coins and rolling dice, of the laws of probability that these activities demonstrate and the statistical predictions that may be derived from them.

En dat is een wat merkwaardige beperking. Munten, dobbelstenen, kansberekening en statistiek hebben allemaal met ‘toeval’ te maken. Maar er zit (soms) ook iets niet-toevalligs in de fout. Nou ja, misschien is het wel *toevallig*, maar *varieert het niet*. Als hij bij elke meting even groot is, kun je dat niet met statistische theorieën over munten en dobbelstenen komen. (Tenzij je dobbelstenen hebt waarmee je altijd ‘zes’ gooit.)

Dat alleen het toevallige deel soms ten onrechte als onzekerheid beschouwd wordt, zal twee oorzaken hebben.

- Iets wat de hele tijd gelijk blijft, komt niet over als onzeker.
- Welk deel fout is én de hele tijd gelijk blijft, wordt niet zichtbaar door de meting te herhalen.

Berendsen (2011) geeft in zijn boek een driedeling die we niet overnemen.

3.1 Classification of errors

There are several types of error in experimental outcomes:

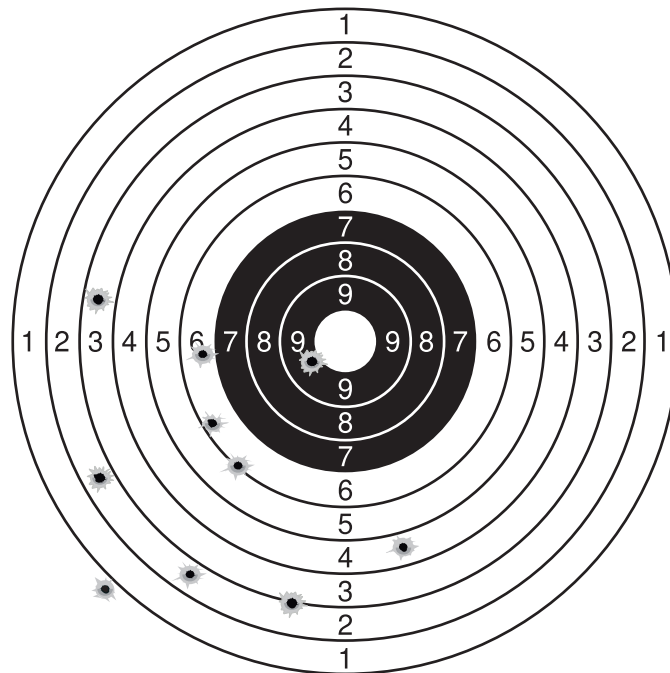
- (i) (accidental, stupid or intended) mistakes
- (ii) systematic deviations
- (iii) random errors or uncertainties

The first type we shall ignore. Accidental mistakes can be avoided by careful checking and double checking. Stupid mistakes are accidental errors that have been overlooked. Intended mistakes (e.g. selecting data that suit your purpose) purposely mislead the reader and belong to the category of scientific crimes.

Je zou *mistakes* inderdaad kunnen opdelen in opzettelijk en per ongeluk. (Niet in *stupid* trouwens, want een fout maken is niet stom. Niet toegeven is dat wel, maar een fout maken is vergefelijk en leerzaam en niet erg. De fout zelf kan wel stom zijn. Degene die hem maakt is dat niet.)

Als mijn personenweegschaal steeds iets anders aangeeft als ik erop en eraf stap, weet ik heus wel dat mijn gewicht (massa) niet met een paar ons verandert in die korte tijd. Dat toevallige deel zie ik wel. Maar dat ik er de hele tijd een procent naast zit, blijkt nergens uit. (Behalve dan dat ik heus wel lichter ben dan wat dat ding zegt... Een paar jaar geleden gaf hij minder aan.)

Toevallige en systematische fouten worden wel eens in kaart gebracht en uitgelegd aan de hand van schietschijven en pijlen die in de roos belanden of juist niet. Die komen we in de loop van dit boek vaker tegen.



Figuur 10.1 Metafoor van toevallige en systematische fouten: de schietschijf

Er zit een spreiding in de kogelgaten, veroorzaakt door de schutter en/of het geweer. Die spreiding duidt op een *toevallige fout*. Dat de meeste kogels in het kwadrant linksonder zitten en dat ze gemiddeld genomen een stuk van de roos af zitten, duidt op een *systematische fout*. Bij de schietschijf kun je beide zien:

- de toevallige fout zit in de grootte van de spreiding
- de systematische zit in de gemiddelde afstand tot de roos

Bij een meting hebben we het probleem dat we de schietschijf niet hebben... We zien wel een spreiding en weten hoe groot die is, gewoon door te kijken naar de kogelgaten (of de metingen). Maar we weten niet hoe ver die afstand tot de roos is.

Hier zit dus de verwarring. Als we een tijd meten en we beweren dat we een systematische fout hebben van 0,2 seconde, dan ben je geneigd om te zeggen: 'Man, corrigeer daar dan voor!' Het probleem is dat je (om een of andere reden) wel weet dat de grootte ongeveer 0,2 seconde is, maar niet of het nou 0,1 s of -0,2 s of 0,15 s of 0,25 s of -0,005 s is. Het zijn in ieder geval geen tien seconden, maar heus wel meer dan een nanoseconde.

We hebben dus een ruwe 'maat' voor die systematische fout. Je weet bijvoorbeeld dat je er grofweg 2% naast zal zitten. Maar je weet niet hoeveel precies. Het is een soort stelpost: je weet dat er nog wat aan mankeert, maar niet precies wat en hoeveel.

We weten van de systematische fout niet eens of hij positief of negatief is. Als we wisten dat de afwijking in de tijd ergens tussen de 0,1 s en 0,3 s zou zitten, zouden we bij alle waarden 0,2 s optellen. Omdat we dan dichterbij komen.

Er is nog iets te zeggen over systematische en toevallige fouten, namelijk dat je beter kunt spreken van een *deterministisch* en een *stochastisch* deel. Een stuk dat de hele tijd hetzelfde is een stuk dat bij elke meting anders kan zijn. Die systematische fout kan door toeval veroorzaakt zijn. Maar dat is het punt niet: als dat toevallige deel de hele tijd hetzelfde is, noemen we het nog steeds systematisch.

Statistisch geformuleerd: als fouten normaal verdeeld zijn met een gemiddelde μ en een standaardafwijking σ , dan is

- de systematische fout gelijk aan μ
- de toevallige fout een normaal verdeeld proces met een gemiddelde van 0 en een standaardafwijking σ

Bij moeilijke woorden zoals *deterministisch* en *stochastisch* kan het leerzaam zijn om te kijken of je er iets eenvoudigers voor kunt vinden. Wie onnodig een vreemde term gebruikt, verbloemt soms dat hij zij het eigenlijk niet snapt.

Soms helpt het al om bij Google het woord in te tikken, vergezeld van 'wiki' of 'Wikipedia'. In dit geval ontdek je dan dat er over *stochastisch* veel meer geschreven wordt dan over *deterministisch*. Dat is ook wel logisch. Een *deterministisch* proces levert telkens hetzelfde op. Daar is weinig méér over te melden. Een *stochastisch* proces is nogal grillig. Telkens komt er iets anders uit, zuiver toevallig. Maar over dat toeval valt wel iets te zeggen! Er (b)lijkt een *kansverdeling* aan het proces ten grondslag te liggen. Normaal gesproken ken je die verdeling niet, maar die is wel te bepalen of te schatten.

Soms helpt het ook om het bij Google in te tikken en het vergezeld te laten gaan van 'etymologie'. Dat is op zichzelf ook een vreemde kreet, die '(studie van de) herkomst van woorden' betekent. Je ontdekt dan:

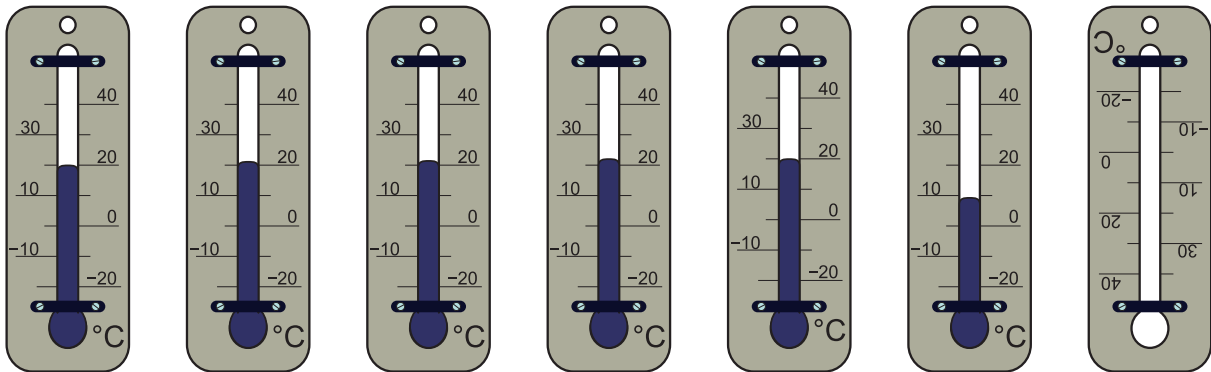
- *Deterministisch* komt van het Latijnse *determinare* (waarin *de* (van) en *terminus* ('eind' of 'grens') zit). Het betekent *bepalen*. De uitkomst is van tevoren *bepaald*. Die ligt dus vast, net als dat systematische deel van de fout. (Systematisch is gewoon een moeilijk woord voor 'stelselmatig'.)
- *Stochastisch* komt van het Griekse *στόχος* dat 'gok' (of doel) betekent. Gokken of gissen doe je als je iets niet weet. Meestal weet je wel *iets*, bijvoorbeeld als je tijdens een kans- of gokspel met een twee dobbelstenen gooit. Het is niet verstandig om dan op 1 of 13 in te zetten.

Over stochastische processen gaat de sectie *Statistiek* in *Metten & Weten*. Als je er écht iets over wilt weten, zijn er betere en uitgebreidere boeken te vinden.

En dan nu het gekste van het hele verhaal. We zullen er telkens tegenaan lopen dat het *deterministische* deel van een fout *niet* te bepalen is en het *stochastische* deel juist *wel*. Als je een meting oneindig vaak zou herhalen – zoveel tijd heb je niet in je aardse bestaan – kon je die kansverdeling exact vaststellen. Van dat *deterministische* deel heb je dan nog steeds geen flauw idee. Je zou dus moeten gokken... of op een andere manier aan kennis moeten zien te komen.

11. Toevallige systematische fouten

Dat de begrippen “toevallig” en “systematisch” soms afhangen van de context, wordt duidelijker als we het volgende voorbeeld bekijken.



Figuur 11.1 Zeven in serie geproduceerde thermometers

We gaan een temperatuur meten met een thermometer. We hebben zeven exemplaren ter beschikking. Hoe zit het met de toevallige en systematische fouten bij het werken ermee?

Er zal een toevallige fout zijn als gevolg van het niet exact kunnen aflezen. Het stukje dat je “ernaast zit”, is willekeurig en daar zit dus spreiding in.

Er zal ook een systematische fout zijn als gevolg van het niet perfect kunnen fabriceren. De hoeveelheid vloeistof in het glazen buisje en de exacte positie van de streepjes van de temperatuurschaal zijn onvolmaakt. Het lukt de fabrikant niet om tot op het molecuul nauwkeurig af te vullen en lijntjes te trekken. Daar zullen procenten aan schommeling in zitten.

Die hoeveelheid vloeistof is dus iets te groot of iets te klein en we weten niet hoeveel. Maar hij is wel de hele tijd gelijk voor één bepaalde thermometer. De gevolgen van de fout zijn dus telkens even groot.

Laten we de zeven thermometers van Figuur 11.1 eens beter bekijken. De eerste vijf (van links) geven allemaal ongeveer 20 °C aan, maar er zitten wel kleine verschillen in, als gevolg van de hoeveelheid vloeistof die niet exact klopt en een bepaalde spreiding heeft. De zesde en zevende thermometer wijken veel af. De zesde staat op ongeveer 10 °C, als gevolg van het feit dat de vloeistof bijna op was toen deze gevuld werd. Bij de zevende was het helemaal op... Er zit dan ook geen druppel in. En tot overmaat van ramp zijn de streepjes ook nog eens op zijn kop afgedrukt!

Het voordeel van strepen die helemaal op hun kop staan, is dat je duidelijk ziet dat er iets misgegaan is bij de productie. Maar wat nu als de streepjes een

klein stukje te laag staan? Niet als gevolg van onvermijdbare onzekerheid, maar als gevolg van een fout bij het produceren? Eigenlijk heb je dan een meetinstrument dat gewoon niet goed is. Terugsturen en een nieuwe vragen. Maar hoe kom je erachter dat hij niet goed is, als het verschil zo klein is?

De twee meest rechtse thermometers kunnen we buiten beschouwing laten. Maar hoe zit het met die linkse vijf? Die zijn gewoon “goed” en hebben een acceptabele spreiding. De specificaties geven ook aan dat de onzekerheid twee graden Celsius is. (Het waren goedkope thermometers, van de Action of AliExpress...)

Nu de essentie van dit verhaal. De hoeveelheid vloeistof verschilt per thermometer en is een gevolg van *toeval*. Maar ook al is de *oorzaak* toevallig, die hoeveelheid vloeistof is voor *elke individuele thermometer wel constant*. Je zou het afvullen als een stochastisch proces kunnen beschouwen, met de hoeveelheid vloeistof in elke thermometer als één concrete uitkomst daarvan.

De afwijking van de ‘juiste’ hoeveelheid vloeistof ligt na de fabricage vast, dus de fout is *systematisch*. Elke keer als je meet met dezelfde thermometer, is de onzekerheid als gevolg van deze afwijking hetzelfde. De situatie wordt anders als je vijf metingen doet en je gebruikt de linkse vijf thermometers alle vijf één keer. Elke meting heeft dan een andere afwijking ten gevolge van die variatie bij de productie. Dan is de fout die onveranderlijk is per thermometer opeens een toevallige fout.

Als de hoeveelheid vloeistof te groot is, zal een te hoge temperatuur worden weergegeven. Het aantal milliliters afwijking is vast, maar de afwijking in de meting als toenemen voor hogere temperaturen. Er zal een lineair verband zijn tussen de gemeten temperatuur en de werkelijke temperatuur (in kelvin!)

Oefeningen

1. Als je deze zeven thermometers ter beschikking hebt, hoe zou je daar dan het beste mee kunnen meten?
2. Als er in een thermometer 5% minder vloeistof zit dan erin zou moeten zitten, is er dan sprake van een onzekerheid of een productiefout?

12. Nauwkeurig & Precies

Als je wilt weten wat een woord betekent, kun je het intikken bij Google of opzoeken in een woordenboek. Dat laatste is wat betrouwbaarder, want via een zoekmachine weet je maar nooit waar je terechtkomt. En als je toch internet hebt, kun je in de gratis online versie van *Van Dale* kijken.

We willen weten wat het verschil is tussen *nauwkeurig* en *precies*.

Volgens *Van Dale* (online):

nauw·keu·rig (bijvoeglijk naamwoord, bijwoord) juist, zorgvuldig, stipt; = accuraat

pre·cies (bijvoeglijk naamwoord, bijwoord; vergrotende trap: *preciezer*, overtreffende trap: *preciest*) zonder afwijking of verschil; juist, nauwkeurig, stipt: *precies honderd euro; hij is al te precies nauwgezet*

Waarom er bij ‘precies’ een vergrotende en overtreffende trap staat en bij ‘nauwkeurig’ niet, is niet duidelijk. Verwarrender is dat bij ‘precies’ het woord ‘nauwkeurig’ als synoniem staat, maar omgekeerd lijkt ‘nauwkeurig’ geen synoniem te zijn van ‘precies’. Hoe zit dat precies?

De tekstverwerker *Microsoft Word* laat je gemakkelijk (door met rechts op een woord te klikken) een synoniem vinden. En of je nu op ‘nauwkeurig’ of ‘precies’ klikt: in beide gevallen verschijnt ‘accuraat’ of ‘nauwgezet’.

Op de website bij het boek *Schrijfwijzer* van Jan Renkema staat ook een toelichting op deze woorden.

De woorden worden soms door elkaar gebruikt, maar verschillen wel in betekenis.

nauwkeurig: zorgvuldig, met aandacht en oog voor details

“Nog nooit werd een luchtweginfectie zo nauwkeurig onderzocht als corona.”

“Ik heb alles nauwkeurig schoongemaakt, je ziet er niets meer van.”

precies: exact, zonder afwijking

“Er komen in elk geval meer dan honderd mensen, maar het precieze aantal is nog niet bekend.”

“Hij begreep niet precies wat zij bedoelde.”

In het woord nauwkeurig zitten de woorden ‘nauw’ voor ‘smal’ en ‘keur’ dat van ‘kiezen’ komt. Dus in nauwkeurig kies je tussen mogelijkheden die dicht bij elkaar liggen.

In het vakgebied van *Metten & Weten* zijn nauwkeurigheid en precisie twee wezenlijk verschillende dingen.

- nauwkeurigheid is de mate waarin het gemiddelde van (oneindig) veel metingen overeenkomt met de werkelijke waarde
anders gezegd: hoe dicht de metingen bij de werkelijke waarde liggen
- precisie is de mate waarin herhaalde metingen met elkaar overeenkomen
anders gezegd: hoe dicht de metingen bij elkaar liggen

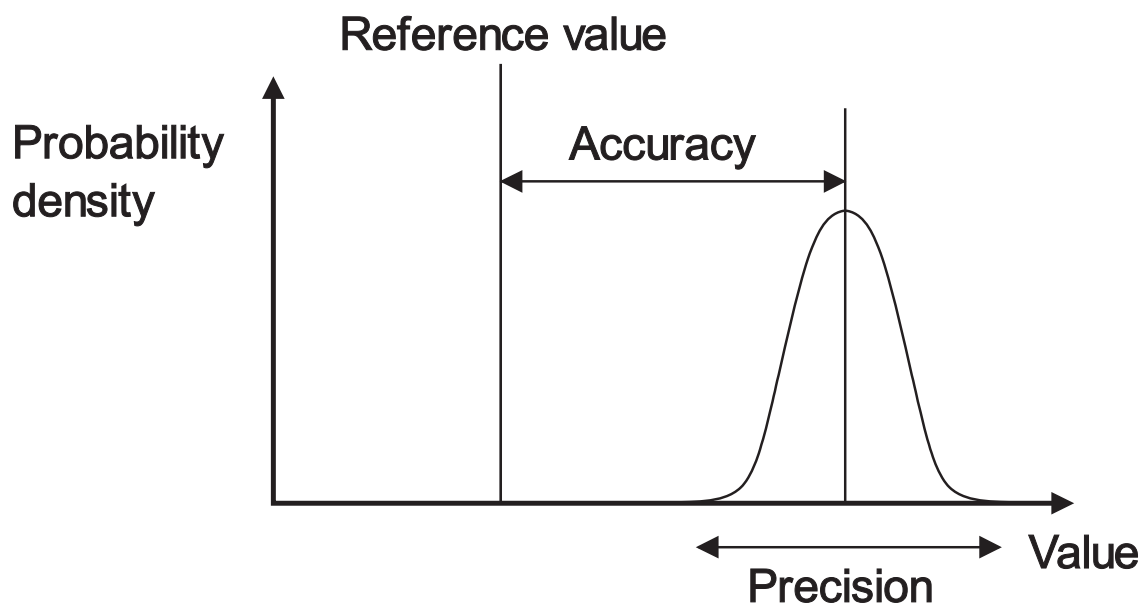
Renkema legt aan de hand van het woordje ‘nauw’ uit dat het gaat *mogelijkheden die dicht bij elkaar liggen*. Dat helpt nog niet echt. Als de metingen *onderling* dicht bij elkaar liggen, noemen we ze *precies*. (De onderlinge spreiding is dan klein.) Als ze gemiddeld dicht bij *de werkelijke waarde* liggen, noemen we ze *nauwkeurig*. (Het verschil tussen werkelijk en gemeten is dan klein.)

Een domme manier om het onderscheid te onthouden, is kijken naar de letters in het woord ‘spreiding’. De s, p, r, e, i en n komen allemaal voor in *precision*. Voor *accuracy* geldt dat geen enkele letter daarvan in ‘spreiding’ zit.



O ja, toch wel: de letter r. Zoals gezegd: het is een domme manier.

De formuleringen zijn soms wat onhandig. Kijk maar naar het plaatje hieronder, dat is overgenomen van de Engelstalige *Wikipedia*.



Figuur 12.1 Het verschil tussen Accuracy en Precision volgens Wikipedia, grafisch weergegeven

De Engelse vertaling van ‘nauwkeurigheid’ is *accuracy* en van ‘precisie’ is het *precision*. Op grond van het plaatje zou je denken dat de nauwkeurigheid toeneemt als de verticale lijnen verder uit elkaar liggen en dat de precisie toeneemt als de klokvormige curve breder is. Dat is juist niet het geval. Het zou beter zijn om te spreken van onnauwkeurigheid en imprecisie. Als je zegt dat je 1,71 meter lang bent en dat vastgesteld hebt met een nauwkeurigheid van 1 cm, bedoel je dan dat je er maar liefst 1,70 meter naast kunt zitten? (“Misschien ben ik wel drie meter veertig, de nauwkeurigheid is nogal klein...”)

Er is nog iets vervelends aan de hand. De termen in de Engelstalige afbeelding (en hun vertalingen in het Nederlands) zijn zeer gangbaar (geweest) en komen ook in veel boeken voor. De officiële termen zijn vastgelegd in ISO 5725. Wat in bovenstaand plaatje *accuracy* is (nauwkeurigheid), heet volgens die norm *trueness* (‘juistheid’). Dat laatste overigens nog met een subtiële toevoeging: de pijl in het plaatje moet korter zijn, omdat we als getalswaarde aan de *precision* de standaardafwijking van de verdeling koppelen. De pijl bij het woord ‘Precision’ in Figuur 12.1 is meer dan vier keer zo groot.

Het wordt nog vervelender: de term *accuracy* (nauwkeurigheid) wordt in ISO 5725-1 gebruikt om te beschrijven hoe dicht een meting bij de werkelijke waarde ligt. Als er dan sprake is van een reeks metingen van dezelfde meetgrootte, gaat het om

- een deel dat wordt veroorzaakt door willekeurige fouten en
- een deel dat wordt veroorzaakt door systematische fouten.

Dan is *juistheid* de mate waarin het gemiddelde van een reeks meetresultaten de feitelijke (ware) waarde benadert en *precisie* de mate van overeenstemming tussen een reeks resultaten.

We begonnen dit hoofdstuk met het opzoeken van *nauwkeurig* en *precies* in *Van Dale* (*online*). En we eindigen met het opzoeken van *exact*.

exact (*bijvoeglijk naamwoord, bijwoord*) nauwkeurig, stipt: exacte wetenschappen wis-, schei- en natuurkunde; *de exacte cijfers heb ik niet; zij heeft exact dezelfde jas als ik* precies dezelfde

Gekker moet het toch niet worden! Als eerste betekenis van ‘exact’ wordt door *Van Dale* ‘nauwkeurig’ gegeven. En ‘exact dezelfde jas’ is volgens datzelfde gezaghebbende boek ‘precies dezelfde’... Pak je een naslagwerk erbij, dan wordt het één pot nat. Ze nemen het niet zo nauw, keurig uitgedrukt. Terwijl dat nou juist precies de bedoeling is!

Daar kunnen de makers van het woordenboek ook niets aan doen. Ze beschrijven hoe gewone Nederlands-sprekende mensen taal gebruiken.

Dat doen we altijd erg slordig. Nou ja, niet altijd. Meestal. (Vaak, in ieder geval. Denk ik.)

Voor de liefhebber: volgens *Van Dale* kun je het tweelettergrepige ‘exact’ niet afbreken aan het eind van de regel. Terwijl dat met het woord ‘exacte’ wel kan.

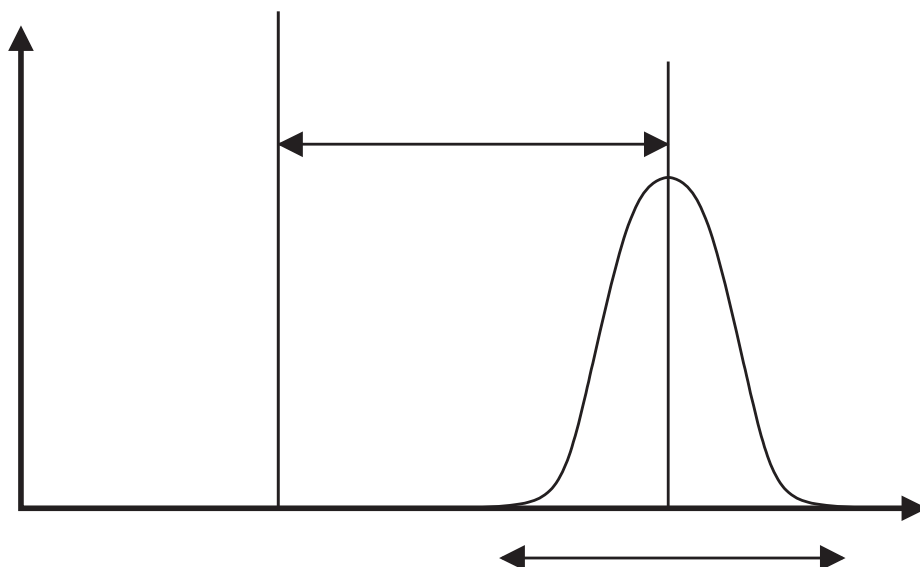
Naast *nauwkeurig* (gemiddeld genomen dicht bij de werkelijke waarde) en *precies* (weinig spreiding in de meetwaarden) hebben we ook nog een derde begrip: *exact*. Een meting is *exact* als de uitkomst in alle mogelijke decimalen overeenkomt met de werkelijke waarde en er geen spreiding is. De onnauwkeurigheid en de imprecisie zijn nul. Er is dan geen verschil tussen meting en werkelijkheid. Bij een telling kan dat het geval zijn.

Voorbeelden van exacte dingen:

- een inch is exact 2,54 cm (want hij is zo gedefinieerd)
- de lichtsnelheid is exact 299 792 458 m/s (want ze is zo gedefinieerd)
- er gaan exact twaalf eieren in een dozijn (in een dozijn *eieren* dan...)

Oefeningen

1. Vul in Figuur 12.2 de juiste Nederlandse woorden in.
2. Vul in Figuur 12.2 de juiste Engelse woorden in.
3. Iemand vraagt naar de lichtsnelheid afgerond op een geheel aantal inches per seconde. Is het mogelijk om een exact antwoord te geven?
4. Maak het woord ‘precision’ met de letters van ‘spreiding’.



Figuur 12.2 soorten onzekerheden, zonder tekst

13. Nauwkeuriger & Preciezer

Een van mijn vroegste jeugdherinneringen – denk ik – gaat terug naar het antwoord dat ik kreeg van mijn vader op een in mijn ogen eenvoudige vraag. “Papa, vijftientig gulden, is dat veel?” Omdat jeugdige lezers geen idee hebben van guldens en florijnen, spieken we even op de [Inflatiecalculator](#) om dit om te rekenen naar euro’s en te corrigeren voor inflatie.

Bedrag 25 Jaartal 1973	Een bedrag van 25,00 gulden uit 1973 is nu € 53,03 op het prijskaartje. Een bedrag van € 25,00 in 2023 was in 1973 evenveel als 11,79 gulden of omgerekend € 5,35.
--	--

Of mijn vraag uit 1973 dateert, weet ik niet. Maar ik schreef dit hoofdstuk in 2023 dus het is wel aardig om precies een halve eeuw eerder te nemen. Bovendien is 1973 het eerste jaar (als je terugtelt in de tijd) waarin het bedrag in euro’s meer dan het dubbele is van het bedrag in guldens.

Een bedrag van € 100,00 in 2023 was in 1973 evenveel als 47,14 gulden of omgerekend € 21,39.

Het leven is dus $€ 100,00 / € 21,39 \approx 467\frac{1}{2} \%$ duurder geworden. Ruim vierenhalf keer zo duur. Als de inflatie in de komende vijftig jaar zo doorgaat, betaal je op een terras in een grote stad in 2073 een bedrag van € 17 voor een cappuccino. Over vijftig jaar... een flink aantal van de huidige studenten zal dat nog meemaken.

Die vraag van mij of vijftientig gulden (nu zou dat zijn: vijftig euro) veel is, werd beantwoord met: “Dat ligt eraan waarvoor.” En dat is ook zo. Maar ik vond destijds dat je niets aan dat antwoord had. Vertel nou gewoon maar of het veel is, zonder daar omheen te draaien.

Van 1973 terug naar *Metten & Weten* en iets wat hierop lijkt. Of iets ‘nauwkeurig’ is of ‘precies’, is namelijk ook niet zomaar te beantwoorden. Je kunt alleen zeggen dat de ene manier van meten *nauwkeuriger* (of *preciezer*) is dan de andere. En je kunt ook nog kiezen of je dat absoluut of relatief neemt.

Hughes & Hase (2010) schrijven iets over die begrippen ‘nauwkeurig’ en ‘precies’. Maar dan in het Engels natuurlijk: *accurate and precise*.

There are two terms that have very different meanings when analysing experimental data. We need to distinguish between an accurate and a precise measurement. A precise measurement is one where the spread of results is 'small', either relative to the average result or in absolute magnitude. An accurate measurement is one in which the results of the experiment are in agreement with the 'accepted' value. Note that the concept of accuracy is only valid in experiments where comparison with a known value (from previous measurements, or a theoretical calculation) is the goal – measuring the speed of light, for example.

Het doet me denken aan die vraag van mijzelf aan mijn vader, ergens in de eerste helft van de jaren zeventig van de vorige eeuw. “Is vijftientig gulden veel?” Ik had ook kunnen vragen: “Is een onzekerheid van 25 cm groot?” Als je bij het bepalen van je eigen lichaamslengte zo’n foutmarge opgeeft, ben je ongeveer net zo lang als iedereen. “Ik ben één meter tachtig, maar ik kan er vijftientig centimeter naast zitten.” Maar als je 25 cm onzekerheid wilt bereiken bij het bepalen van de afstand tussen de aarde en de maan, dan wordt dat een ander verhaal. Op een lengte van 1,80 meter is 25 cm bijna 14%. Op 380.000 km is het minder dan zeven miljoenste van een procent.

Nauwkeurig en precies zijn dus relatieve begrippen. En wat nu zo bijzonder is: ook de *relatieve* onzekerheden zijn getallen die je moet relateren aan iets anders. “Papa, is een relatieve onzekerheid van 5% klein?” “Dat ligt eraan...” Ook een relatieve fout is betrekkelijk.

Eigenlijk kun je niet zeggen dat een meting nauwkeurig of precies is. Je kunt wel zeggen dat de ene meting preciezer (of nauwkeuriger) is dan de andere. Of de nauwkeurigheid en precisie onderling met elkaar vergelijken.

Oefeningen

1. Twee multimeters worden gebruikt voor het bepalen van een spanning. De eerste heeft als meetresultaat $2,36 \pm 0,02$ V.
De tweede heeft als meetresultaat 2350 ± 5 mV.
Welke van beide multimeters is nauwkeuriger?
2. Twee digitale afstandsmeters worden gebruikt voor het bepalen van een lengte. De eerste heeft een spreiding van 2% op een lengte van ongeveer acht meter. De tweede heeft een spreiding van 1 cm op een lengte van ongeveer zestig centimeter.
Welke van beide afstandsmeters heeft de grootste nauwkeurigheid?

14. Spreiding in ruimte en tijd

In hoofdstuk 1 verscheen het *definitieprobleem*. Wat bedoel je precies als je het hebt over de hoogte van een deur? De maximale hoogte? De gemiddelde hoogte? Meestal zit er in praktijk niet veel verschil in. Maar als je in een huis van tweehonderd jaar oud woont, weet je dat het ook anders kan zijn.

Wat bedoel je precies als je vraagt naar de temperatuur in een kamer? Veel mensen kijken op een thermometer die aan de muur hangt en noemen het aantal graden Celsius. (Meestal laten ze het woord Celsius weg... Dat mag, zolang je het maar niet doet op je werk of als student op een tentamen, of tijdens de les.)

Dat die graden heus niet in Fahrenheit zijn, dat nemen anderen wel van je aan. Een belangrijkere vraag is: wat meet je precies? Als je een thermometer eerst op de grond zet en daarna boven op een kast, zul je zien dat je niet precies dezelfde temperatuur vindt. Dat heeft te maken met het feit dat warme lucht stijgt en koude lucht daalt. En misschien ook wel met een verandering in de temperatuur in de kamer.

De temperatuur in een kamer varieert in de ruimte en in de tijd.

Wiskundig gezien zou je het antwoord op de vraag ‘Wat is de temperatuur in deze kamer?’ alleen maar kunnen beantwoorden in de vorm van een functie die afhangt van de plaats (drie coördinaten) en de tijd (één variabele).

$$T(x, y, z, t)$$

Onder ideale omstandigheden zou dat kunnen. Maar dan ontstaat het volgende probleem. In hoeveel decimalen geef je de coördinaten weer? En in hoeveel cijfers achter de komma noteer je de tijd? Niet te vergeten: hoeveel cijfers gebruik je voor de temperatuur zelf?

Zelfs als je in micrometers en nanoseconden de temperatuur in de ruimte zou kunnen bepalen en je voor de temperatuur vier of tien of honderdzestig cijfers achter de komma zet: er komt ergens een keer een moment dat het niet exact meer is. We hebben bij volmaakte meetinstrumenten dus al te maken met twee soorten onzekerheid als gevolg van datgene wat we meten:

1. Alles gaat op één hoop, terwijl er een spreiding in ruimte en tijd is.
2. Er is een eindig aantal cijfers bij het noteren van de uitkomst.

Algemener gezegd:

1. Er zit een onzekerheid in datgene *wat* we meten.
2. Er zit een onzekerheid in datgene *waarmee* we meten.

Vaak gaat alle aandacht uit naar die tweede vorm van onzekerheid: de beperkingen van je meetinstrument en het aflezen daarvan. Daar heb je ook nog eens twee soorten in:

- 2a. Systematische fouten, die altijd hetzelfde zijn.
- 2b. Toevallige fouten, die per meting verschillen, maar gemiddeld wel nul zijn.

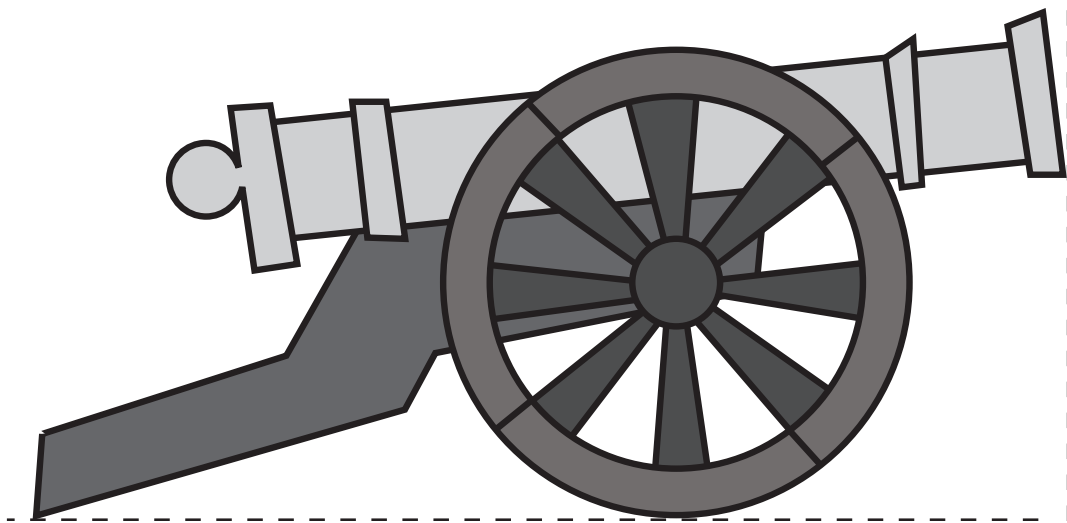
Veel lesboeken op dit vakgebied gaan uitgebreid in op punt 2 en het verschil tussen 2a en 2b, maar veel minder op punt 1. Helaas wordt er als het om onzekerheden gaat onevenredig veel aandacht besteed aan 2b. Want hoewel toevallige fouten ‘lastiger’ lijken, zijn ze dat juist niet. Ze vallen veel meer op dan systematische. Bovendien kun je er ook iets mee: statistiek bedrijven en ze daarmee verkleinen.

15. Spreiding in wat we meten

“Hoe ver kun je schieten met een kanon?”

Hebben we bij deze vraag een definitieprobleem? Ja. Op z’n minst eentje: welk kanon bedoel je? Het ene kanon is het andere niet. Laten we de vraag aanpassen.

“Hoe ver kun je schieten met dit ene kanon uit Figuur 15.1?”



Figuur 15.1 Kanon

Natuurlijk kun je met een getekend kanon niet schieten, maar dit is ook maar een boek. Laten we uitgaan van één specifiek werkelijk bestaand kanon.

Hebben we nu nog een definitieprobleem? Je zou in ieder geval goed moeten vastleggen welke afstand je precies bedoelt. Laten we meten vanaf het kruispunt van de stippelijnen: het uiteinde van de loop van het kanon en de grond zijn dan je uitgangspunten. En we trekken loodrecht een lijn omlaag.

Zijn die lijnen in de werkelijkheid exact gedefinieerd? De horizontale lijn is wat twijfelachtig: de ondergrond is namelijk nooit volmaakt vlak. Maar je zou de verticale lijn wel kunnen definiëren als de exacte richting van de zwaartekracht. Niet dat je die exact kunt *meten*, maar je hebt dan wel een exacte *definitie*.

Missen we iets? We hebben een werkelijkheid met drie ruimtelijke dimensies. (Misschien zijn er ook hogere dimensies, maar die zien we niet en kunnen we dus bij het meten buiten beschouwing laten.)

Toch hebben we nu aan één punt genoeg om het andere te vinden.

- vanaf het exacte midden van de loop van het kanon
- in de richting van de zwaartekracht
- naar de grond toe

We gaan voor het gemak – en voor onze veiligheid – uit van een speelgoedkanon dat een paar meter kan schieten. Dan kunnen we onze meting doen met een rolmaat van 5 meter. Hoe groot is de onzekerheid in de meting? Een paar millimeter? Een hele centimeter? Daar kun je verschillend over denken. En dus is het erg belangrijk dát je er in praktijk over nadenkt.

Dit zijn de afstanden die we meten:

2,60 m, 2,73 m, 2,57 m, 2,69 m, 2,56 m, 2,68 m, 2,72 m, 2,69 m, 2,60 m, 2,67 m.

Het verschil tussen de grootste gemeten afstand (2,73 m) en de kleinste gemeten afstand (2,56 m) is 17 cm. Dat is heus wel meer dan die onzekerheid in het aflezen van het meetlint.

We zien variaties in verschillende soorten:

1. De exacte definitie.
2. De kwaliteit van je meetlint.
3. Het aflezen van het meetlint.
4. De variatie in het schieten.

Het eerste punt valt hier op te lossen door eenduidig te definiëren. Het tweede en derde punt vormen de onzekerheid in je meting (*measurement*). Het vierde punt is de onzekerheid in datgene wat je meet (*measurand*).

De variaties hebben te maken met luchtweerstand en toevallige windkracht, hoeveelheid kruit, positionering van de kogel en nog wel wat dingen. Die spreiding, is zo vreemd niet.

We hebben hier dus te maken met een spreiding in datgene wat je meet die groter is dan je onzekerheid in de meting.

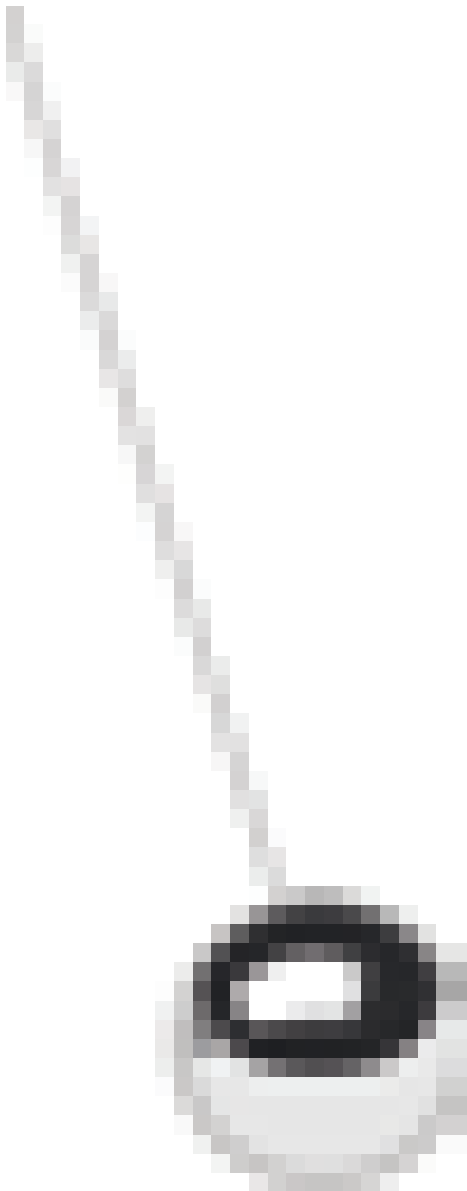
Oefeningen

1. Geef aan hoe je de onzekerheid in de meting zou kunnen schatten.
2. Geef een aanduiding voor de spreiding in de meting zelf.

16. Spreiding door het meten

De spreiding in de afstanden die we gemeten hebben bij dat speelgoedkanon uit het vorige hoofdstuk werd vooral veroorzaakt door dat kanon en de kogels zelf. Een beetje door de wind misschien. Het was wel telkens hetzelfde kanon, maar niet altijd dezelfde kogel.

Het meetlint had een onzekerheid bij het aflezen en ook het lint zelf is niet perfect. Zoals de Engelsen zeggen: *Nobody's perfect, even not a meetlint.*



Figuur 16.1 Perfecte kogel aan een perfect koordje

Laten we nu eens gaan kijken naar één zo'n kanonskogel. Een heel erg mooie kogel, mooi zwaar en symmetrisch. Opgehangen aan een flinterdun koordje.

We gaan nu de slingertijd meten van deze bijna volmaakte bol, in een ruimte waar alle ramen dicht zijn en de temperatuur en luchtdruk en ventilatie perfect constant is.

De tijd meten we op met een stopwatch. Een ouderwetse stopwatch met knoppen en wijzers, of een moderne app op een telefoon. In beide gevallen zul je zien dat je echt niet altijd precies dezelfde periodetijd meet.

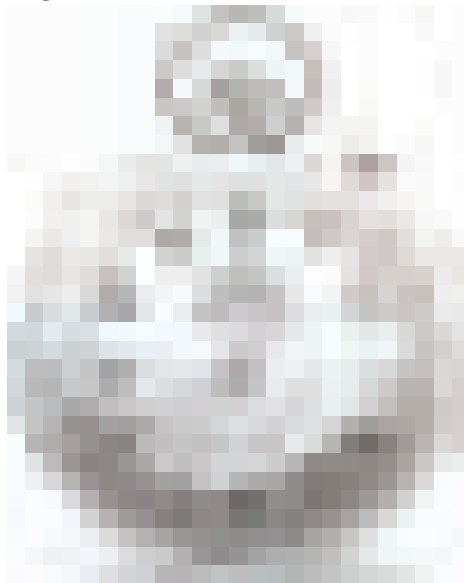
Waar komen die verschillen vandaan?

Het is telkens dezelfde kanonskogel aan datzelfde koordje. Dat zal niet veel slijten of veranderen na één keer slingeren. Ook de hoek is nauwelijks verschillend. We houden de kogel strak tegen de tekst van dit hoofdstuk aan. Kijk maar in Figuur 16.1. Hij raakt nét niet de woorden op de volgende twee regels. Die hoek is telkens perfect gelijk aan de vorige keer. Nou ja, bijna perfect dan. *Nobody's perfect, even not Figuur 16.1.*

Hoe groot is de invloed van de tijdsmeting op de gemeten tijd?

Bij het kanon was het vooral het kanon en de kogel en de wind en het kruit of het onkruid

van het gras waarin de kogel landde. Dat meetlint trof geen blaam, dat kon je goed aflezen.



Figuur 16.2 Een stopwatch, met wijzers

Kun je een stopwatch goed aflezen? De wijzers wel en een digitale stopwatch kun je zelfs *perfect* aflezen. (Als je nog nooit een stopwatch met wijzers gezien hebt, er staat er eentje in Figuur 16.2.) Toch meet je niet elke keer hetzelfde. Dat heeft te maken met hoe goed en snel wij reageren en waarnemen. Onze ogen, hersenen, zenuwen, spieren en vingers zorgen voor de spreiding. We hebben dus twee soorten spreidingen.

1. Spreidingen in datgene wat we meten.
2. Spreidingen door het meten,
 - a. veroorzaakt door het meetapparaat, en
 - b. veroorzaakt door de menselijke waarneming.

Nobody's perfect. Wat je meet is niet perfect, je meter is het niet én jijzelf bent het niet. Dat geeft niet, maar hou er rekening mee en schat de gevolgen ervan goed in.

Oefeningen

1. Als de kogel in Figuur 16.1 zo goed mogelijk (perfect lukt toch niet) tegen de tekst aan wordt gehouden, hoe groot zal dan de spreiding / onzekerheid zijn in de hoek?
2. Als je een reactietijd hebt van 0,3 seconde, hoe groot is dan de onzekerheid als gevolg hiervan in een tijdmeting met een stopwatch?

17. Beïnvloeden van wat we meten

Vaak meten we ‘iets’ en is dat iets niet helemaal scherp gedefinieerd. De hoogte van een deur varieert, dus je zult moeten aangeven of je maximale, gemiddelde of minimale hoogte bedoelt. Dat de hoogte varieert met de temperatuur en luchtvochtigheid, is nog zo’n probleem. We noemden dit eerder al een *problem of definition*: een definitieprobleem.

Je kunt ervoor kiezen om het scherper te definiëren en al de genoemde factoren vast te leggen. Een alternatief is dat je het ‘gewoon’ over de hoogte hebt en die andere zaken meeneemt in je onzekerheid. Als er een spreiding van 3 mm is over het verloop van de deur en die ook nog eens 2 mm kan uitzetten of krimpen door temperatuurschommelingen en vocht, kun je dit natuurlijk ook meenemen in de onzekerheid zelf. De bijdrage van dit deel van de onzekerheid is dan groter dan die van het meetinstrument.

Onderschatting van fouten wordt vaak hierdoor veroorzaakt. We zien dat meetlint en denken: dit is wel vrij goed af te lezen. Maar er is veel meer onzeker dan dat. Dat is wel zeker.

In het genoemde voorbeeld lag het aan de deur zelf. Vaak ‘doen’ metingen ook iets met datgene wat we meten.

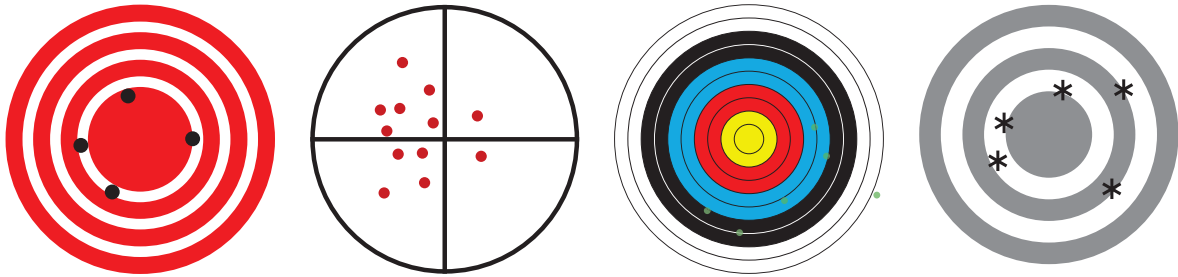
- Een schuifmaat waarmee we een zacht voorwerp opmeten drukt dat zachte voorwerp een klein beetje samen.
- De spanning van een batterij zakt in als er stroom loopt en er loopt inderdaad stroom als je de spanning meet.
- De temperatuur van een vloeistof daalt of stijgt een klein beetje doordat de temperatuur van de thermometer die je erin stopt niet gelijk is aan de temperatuur van de vloeistof.

Ook hier geldt dat je de beïnvloeding zou moeten meenemen in het resultaat. In sommige gevallen kun je ervoor compenseren of rekening houden met het feit dat je weet in welke ‘richting’ de invloed is. Als de schuifmaat inderdaad het voorwerp een beetje indrukt, weet je dat de afgelezen waarde te klein is.

Als je weet dat je moet compenseren maar niet exact met hoeveel, zul je het moeten schatten. Die onzekerheid introduceert een nieuwe toevallige fout.

18. Schietschijven en rozen

Wie met Google het internet afzoekt op ‘accuracy and precision’, vindt een heleboel plaatjes, die er pakweg zo uitzien.



Figuur 18.1 Voorbeelden van schietschijvenplaatjes

Minder mooi dan die schietschijf met kogelgaatjes in dit boek – al zeg ik het zelf – maar wel allemaal met hetzelfde idee. Ronde vormen, meestal met ringen, en daarop verspreide bolletjes, kruisjes, pijlen of punten.

Meestal hebben ze vier van die plaatjes. Die staan dan met z'n vieren naast elkaar (lekker makkelijk) of in een matrix van twee bij twee (overzichtelijk).

Die matrix ligt het meest voor de hand: je hebt het namelijk over twee dingen (nauwkeurig en precies) en een meting kan elk van beide eigenschappen wel of niet hebben.

- A. nauwkeurig en precies
- B. nauwkeurig, maar niet precies
- C. niet nauwkeurig, maar wel precies
- D. niet nauwkeurig, en ook al niet precies

Toch is het leven zo simpel niet, zagen we in het vorige hoofdstuk al. Waar ligt precies de grens tussen precies en niet precies? Hoe nauwkeurig kun je vaststellen of iets nauwkeurig is?

In de plaatjes heb je daar meestal geen last van. Het verschil is vaak volkomen duidelijk. De makers van die schietschijven en rozen tekenen zelf die gaten en punten en er is nauwelijks discussie of er sprake is van een A, B, C of D. Twijfelgevalletjes zitten er niet tussen.

Nu zijn die afbeeldingen natuurlijk niet bedoeld om de werkelijkheid uit te beelden, maar om iets *uit te leggen*. Maar wat ze mij eens zouden moeten uitleggen, is waarom er telkens gebruik gemaakt wordt van ronde, tweedimensionale schietschijven. Meten doe je meestal aan een *scalaire* grootte, niet

aan een vector. En als het wel een vector is, ben je meestal maar in één richting geïnteresseerd, of meet je de richtingen apart op.

Neem bijvoorbeeld die tweede schietschijf van rechts. Eigenlijk ben je alleen geïnteresseerd in de afstand tot het midden (de roos) en de kleur van het gebied waarin hij ligt. Een platte, eendimensionale balk zou volstaan.



Figuur 18.2 Platgeslagen versie van het schietschijvenplaatje

Toegegeven, dat plaatje is minder leuk. Maar dit is wat het is. Eén bolletje ligt in het middelste gebied en twee in de hokjes erbuiten en nog twee in de volgende hokjes. In de zwarte zit niks, in de witte nog eentje. Afwijkingen van de werkelijke waarde (*true value*) kunnen positief of negatief zijn. Maar het is eendimensionaal.

Het plaatje klopt niet met het vorige plaatje, zag je het? Deze balk is een platte versie van die schietschijf, gemaakt door alleen de ‘horizontale coördinaten’ te gebruiken. In feite moet je natuurlijk de afstand vinden tot de roos. Maar het gaat om het idee dat je maar één dimensie hebt.

Hoewel die uitleg met die schietschijven best handig en verhelderend is, gaat die aan één belangrijk aspect nogal eens voorbij. Bij metingen ben je geïnteresseerd in de roos. Waar is die? Dat probeer je vast te stellen door te meten. Maar als je niet weet waar de roos is, waar mik je dan op?

Eigenlijk komt onderstaand plaatje beter overeen met de werkelijkheid.



Figuur 18.3 Schietschijvenplaatje zonder schietschijven

Probleem: we zien nu alleen de bollen, ballen en kogels. Liggen die ver uit elkaar? Tsja... Wat is ver? In het ene plaatje verder dan in het andere. Liggen ze dicht bij de roos? Geen idee... We weten niet eens waar die schijf hangt en eindigt, laat staan waar de roos zit.

19. Onzekerheidsrelatie van Heisenberg

Naast de in het vorige hoofdstuk genoemde vormen van onzekerheid is er in de natuurkunde nog eentje die het vermelden waard is, maar niet gaat over de onzekerheid die in dit boek vooral de aandacht krijgt.

De onzekerheidsrelatie van Heisenberg uit 1927 is een belangrijk onderwerp in de kwantummechanica. Het zegt dat er paren van grootheden bestaan waarvoor geldt dat ze niet tegelijkertijd exact vastgelegd kunnen worden. Plaats en impuls zijn geen scherp gedefinieerde waarden, maar toevalsvariabelen met kansverdelingen, die een gemiddelde en een standaardafwijking hebben. Als de standaardafwijking van de ene kleiner wordt, wordt die van de andere groter. Ze kunnen nooit allebei gelijk worden aan nul.

Voor de ondergrens van de standaardafwijkingen van de kansverdeling geldt

$$\Delta x \Delta p \geq \frac{h}{4\pi}$$

waarin Δx de standaardafwijking (en dus de onzekerheid) in de plaats is, Δp de onzekerheid in de impuls en h de constante van Planck. De waarde van deze constante is gedefinieerd als exact $6,626\,070\,15 \times 10^{-34}$ J.

Laten we voor het gemak eens zeggen dat Δx en Δp waarden hebben die rond de 10^{-17} m en 10^{-17} kgm/s liggen. Dan heb je het niet over micrometers of nanometers of femtometers, maar honderdsten van een femtometer. In de dagelijkse praktijk van de ‘gewone’ metingen die wij doen in een fabriek of laboratorium heeft dat geen betekenis. Het is een onzekerheid die heel klein is, maar ook wel heel interessant. Het gaat hier niet om een onzekerheid in de meting en ook niet om een soort definitieprobleem. We kunnen meten en definiëren wat we willen: dit is iets wat gewoon in de natuur zit. Helaas gaat *Meten & Weten* hier niet over.

De Heisenberg in het plaatje hiernaast is een andere dan die Duitse wetenschapper. Een student die een proefversie van dit boek las, vond het leuk om toch dit portret erin te zetten.

Ik ook. Daarom staat het er nu.



20. Error & Uncertainty

Stephanie Bell (1999) beschrijft in haar boek *A Beginner's Guide to Uncertainty of Measurement* glashelder op wat het verschil is tussen een fout (error) en onzekerheid (uncertainty). Dat het verschil belangrijk (important) genoemd wordt, is terecht.

2.3 Error versus uncertainty

It is important not to confuse the terms 'error' and 'uncertainty'.

Error is the difference between the measured value and the 'true value' of the thing being measured.

Uncertainty is a quantification of the doubt about the measurement result.

Whenever possible we try to correct for any known errors: for example, by applying corrections from calibration certificates. But any error whose value we do not know is a source of uncertainty.

In het Nederlands vertaald door Google Translate:

2.3 Fout versus onzekerheid

Het is belangrijk om de termen 'fout' en 'onzekerheid' niet te verwarren.

Fout is het verschil tussen de gemeten waarde en de 'echte waarde' van het ding dat wordt gemeten.

Onzekerheid is een kwantificering van de twijfel over het meetresultaat.

Waar mogelijk proberen wij bekende fouten te corrigeren: bijvoorbeeld door correcties uit kalibratiecertificaten toe te passen. Maar elke fout waarvan we de waarde niet kennen, is een bron van onzekerheid.

Een fout is wel een bron van onzekerheid, maar het is niet de onzekerheid zelf. De fout kan heel klein zijn terwijl de onzekerheid groot is, als een toevallige fout toevallig klein is. Een fout kan trouwens ook groter zijn dan de onzekerheid. Bij een normaal verdeelde toevallige fout zit 68% van de metingen minder dan één keer de standaardafwijking van het gemiddelde. Bijna één op de drie metingen zit er dus buiten.

We hebben nu twee mogelijke misverstanden.

- Een fout is niet per se ‘iets verkeerd gedaan hebben’. Soms (meestal) ontkom je er niet aan dat er een verschil zit tussen de werkelijke waarde en je meting.
- Een fout geeft aan hoeveel je ernaast zit, een onzekerheid is daar slechts een indicatie van. De eerste ken je niet, de tweede is een schatting. De eerste is een getal dat per meting verschilt, de tweede is een kansvariabele.

Eigenlijk hebben we nu geen *twee* mogelijke misverstanden, maar *drie*. Het derde is het misverstand dat we nu geen twee mogelijke misverstanden zouden hebben, eigenlijk.

Wat Stephanie Bell er niet bij zegt, is dat het vaak door elkaar wordt gehaald. Hughes & Hase benoemen dat wel.

You should note that despite many attempts to standardise the notation, the words ‘error’ and ‘uncertainty’ are often used interchangeably in this context—this is not ideal — but you have to get used to it!

Tsja... Je zult eraan moeten wennen dat veel mensen (en boeken) het verkeerd doen. Maar het is niet handig om daar zelf aan mee te gaan doen of eroverheen te lezen als het niet klopt.

In het Nederlands vertaald door Google Translate:

Houd er rekening mee dat ondanks vele pogingen om de notatie te standaardiseren, de woorden ‘fout’ en ‘onzekerheid’ in deze context vaak door elkaar worden gebruikt – dit is niet ideaal – maar je moet er wel aan wennen!

Wen er inderdaad maar aan. Maar blijf wel hoofdschuddend inzien dat het echt iets anders is.

21. Werkelijk & Geaccepteerd

Dit hoofdstuk is gedeeltelijk een herhaling van wat we al gezien hebben. Het gaat over iets wat vaak verwarring geeft. En die verwarring wordt een beetje in stand gehouden door de term *reference value* in het plaatje van *Wikipedia*.

Wat is nou weer een *reference value* (referentiewaarde)? Is dat dan nóg iets anders dan *true value* (werkelijke waarde) en *accepted value* (geaccepteerde waarde)?

Nee. Als het goed is, bedoelen de makers van het plaatje de *true value*. Maar die kennen we niet. Eigenlijk is dat een tekst om op een spandoek te drukken en dan heel groot op te hangen.

De werkelijke waarde ken je niet.

You do not know the true value.

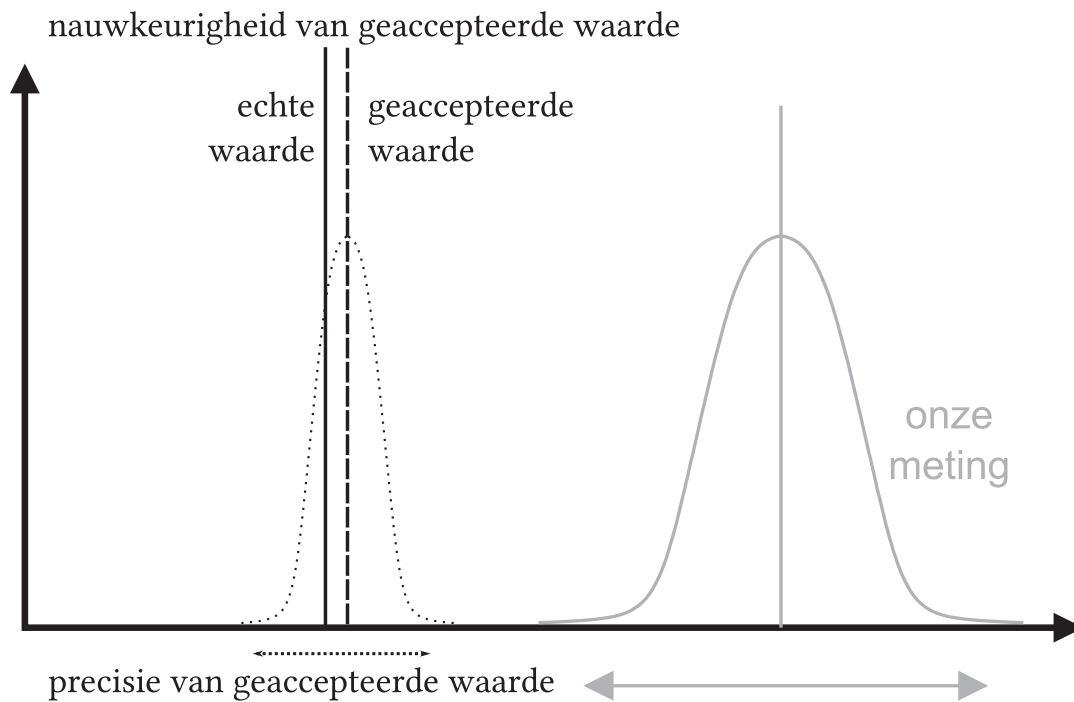
Vreemd eigenlijk, dat je in het Engels het lijdend voorwerp niet zo mooi vooraan in de zin kunt zetten.

De werkelijke waarde ken je niet. Anders zou je ook niet hoeven te meten. Wie gaat nou iets meten wat hij al weet?

Wie de versnelling van de zwaartekracht op zijn plekje op aarde wil weten, kan die meten of opzoeken. De werkelijke waarde ken je niet. Wat je opzoekt, is de *accepted value*. Een of ander genormaliseerd instituut heeft metingen en berekeningen gedaan. Die mensen hebben dat – nemen we aan – héél goed gedaan. Onze eigen uitkomsten leggen we daarnaast en we zien in het verschil onze gebrekkigheid. Als we een strijdigheid vinden, concluderen we dat we zelf iets verkeerd deden. Of dat we te optimistisch waren over de onzekerheid.

De werkelijke waarde ken je niet, en omdat je nu eenmaal behoefte hebt aan iets om mee te vergelijken of rekenen nemen we dan maar iets uit een boek. Vaak is dat een *gemeten* waarde met een (kleine) onzekerheid. Die onzekerheid zal veel kleiner zijn dan onze huis-tuin-en-keukenmeting met een stukje touw dat we in de keukenla vonden en ‘iets zwaars’ om mee te slingeren.

Wat nu als wel de *accepted value* en de *true value* in dat *Wikipedia*-plaatje willen tonen? Dan krijg je Figuur 21.1.



Figuur 21.1 echte en geaccepteerde waarde met onzekerheden

De werkelijke waarde ken je niet. Maar hij bestaat wel! Althans, onder bepaalde omstandigheden. Die valversnelling in *precies* het zwaartepunt van je slinger op die ene plek in het universum, die heeft een waarde die hooguit nog afhangt van hoe de massa van de aarde varieert. Dat zijn heus geen grote veranderingen. Voor een natuurconstante geldt nog sterker: het *is* een werkelijke waarde. Maar die ken je niet.

De geaccepteerde waarde komt dicht in de buurt, maar heeft wel een onzekerheid. Een onzekerheid die is opgebouwd uit een toevallig deel en een systematisch deel.

De verticale dichte lijn in Figuur 21.1 is de werkelijke waarde.

De gestreepte dichte lijn in Figuur 21.1 is de geaccepteerde waarde.

De gestippelde kromme is de kansverdeling van de meetwaarden als je de meting van die geaccepteerde waarde overdeed onder precies dezelfde omstandigheden. Die verdeling is smaller dan “de onze” en het gemiddelde ligt dichter bij de werkelijke waarde. De metingen waren namelijk nauwkeuriger én preciezer. In dat onzekerheidsgebied ligt ergens de werkelijke waarde.

22. Beste schatting van de waarde

Wat is de beste schatting van de werkelijke waarde (*true value*) van een grootheid die je meet?

Als je één meting hebt, is het antwoord niet zo moeilijk: die ene meting. Maar wat als je er twee hebt? Of drie? Of tien, of een paar honderd?

In rekenkunde en statistiek zijn het gemiddelde, de mediaan en de modus gangbare maten. Als we een schatting willen hebben, welke van deze drie is dan de beste?

Als we één spanningsmeting doen en die levert 5 V op, dan is 5 V de beste schatting. Als we daarna nog een meting doen en 7 V vinden, dan wordt het lastiger... Het kan natuurlijk zijn dat de spanning toegenomen is. Dan kun je het beste de laatste waarde nemen. Maar als er geen reden is om dat aan te nemen, is 6 V een aardige keuze.

Hoewel... Als ik 5 V en 7 V meet, en ik neem 6 V als beste schatting, dan kies ik dus voor een waarde die niet gemeten is. Stel nu eens dat je vaker meet, en je komt op de volgende metingen. 5 V, 7 V, 5 V, 5 V, 5 V, 5 V, 5 V, 7 V, 5 V, 5 V, 5 V, 5 V, 5 V, 5 V, 7 V, 5 V, 5 V, 5 V, 5 V, 5 V, 7 V, 5 V, 5 V.

Je meet 21 keer 5 V en 4 keer 7 V. Zou je in dat geval niet gaan denken dat 5 V de “echte” waarde is en 7 V een toevallige verstoring in de meting?

In bovenstaande reeks is het gemiddelde 5,32 V. De modus en de mediaan zijn beide gelijk aan 5 V.

Bij twee metingen (5 V en 7 V) is 6 V ondanks alles een aardige schatting. Het is zowel het gemiddelde als de mediaan. Maar wat doe je als je 5 V, 7 V en 5 V meet? Of 5 V, 7 V, 5 V en 5 V? In het laatste geval van vier metingen heb je een modus en mediaan van 5 V en een gemiddelde van 5,5 V.

Barford (1985) stelt in hoofdstuk 2 van zijn boek dat de modus niet uniek gedefinieerd is. Althans niet in alle gevallen. Als je twee keer 5 V meet en twee keer 7 V, dan zijn er twee waarden die het vaakst voorkomen.

The mean, however, is always uniquely defined and this is another reason, in the absence of arguments to the contrary, for using this and defining the best estimate of a single number physical quantity, derived from n measurements of equal reliability $x_1, x_1, \dots, x_n$ as the mean,

$$X_n = (x_1 + x_1 + \dots + x_n)/n$$

Dat is een zinsnede om over na te denken.

... in the absence of arguments to the contrary...

Zijn er geen argumenten tegen het gemiddelde dan? In het voorbeeld van die spanning die meestal 5 V was en soms 7 V hadden we best wel een reden om anders te kiezen. Wat zouden we doen als we 24 keer 5 V gemeten hadden en één keer 30 V. Het gemiddelde is dan exact 6 V. Toch ligt het meer voor de hand om te denken dat er bij die ene meting iets misgegaan is.

In veel gevallen beschouwen we het gemiddelde van een meting als de beste schatting.

23. Beste schatting van de precisie

Een meting bestaat uit een beste schatting en een onzekerheid, zeggen we vaak. En dat is ook zo. Maar voor je het weet, verlies je uit het oog dat je dan met *twee* schattingen te maken hebt. Eigenlijk zelfs met *drie*.

- een beste schatting van de waarde (schatting 1)
- een beste schatting van de onzekerheid, bestaande uit
 - een deel dat altijd gelijk is (schatting 2)
 - een deel dat varieert (schatting 3)

Van de eerste weten we heus wel dat er een onzekerheid op zit. Daar gaat dit vakgebied zo'n beetje helemaal over. Over die tweede weten we *helemaal niets* op basis van de metingen.

Zoals we er bijna blind van uitgaan dat *het gemiddelde* de beste schatting is van de (werkelijke) waarde, nemen we aan dat de standaardafwijking de beste maat is van de spreiding. Een aanname die voor normale verdelingen zo gek nog niet is.

Barford (1985) laat op een aardige manier zien hoe we tot de gebruikelijke formule voor het schatten van de standaardafwijking komen.

Stel, je hebt één meting. Het is alleszins redelijk om het gemiddelde van die ene meting (de meting zelf dus!) als schatting te nemen voor de waarde. Maar het is *volstrekte onzin* om de spreiding in de meting(en) te nemen als schatting van de onzekerheid. Die ene meting is exact gelijk aan het gemiddelde... de spreiding is dus nul.

No one number can ever give us two pieces of independent information. One measurement gives no information whatsoever about the precision of the apparatus.

Stel, je hebt twee metingen. Het gemiddelde is de som van beide gedeeld door twee. Nu kunnen we wél een zinnige uitspraak doen over de spreiding.

$$\delta_1 = |x_1 - \bar{x}|$$

$$\delta_2 = |x_2 - \bar{x}|$$

Het lijkt nu we te maken hebben met twee gegevens van de spreiding, maar dat is niet zo. Kijk maar naar de formule voor het gemiddelde.

$$\bar{x} = \frac{x_1 + x_2}{2}$$

Als je die formule invult in de twee voor de spreiding vind je:

$$\delta_1 = |x_1 - \bar{x}| = \left| x_1 - \frac{x_1 + x_2}{2} \right| = \left| \frac{2x_1 - x_1 - x_2}{2} \right| = \left| \frac{x_1 - x_2}{2} \right|$$

$$\delta_2 = |x_2 - \bar{x}| = \left| x_2 - \frac{x_1 + x_2}{2} \right| = \left| \frac{2x_2 - x_1 - x_2}{2} \right| = \left| \frac{x_2 - x_1}{2} \right|$$

Met andere woorden: $\delta_1 = |-\delta_2| = \delta_2$

Voor wie niet zo veel met wiskunde heeft, kan het een stuk makkelijker. Als ik 5 V en 7 V heb gemeten, dan is het gemiddelde 6 V. Die 5 V en 7 V zitten beide 1 V van het gemiddelde af. Zo werkt dat nu eenmaal met het gemiddelde van twee getallen: dat zit precies in het midden, en dus even ver van allebei.

We zien nu al één belangrijk ding: we hebben geen twee waarden die aangeven hoe groot de spreiding is, maar slechts eentje.

Wat nu als je de gemiddelde spreiding van N metingen wilt hebben? Elk van die N metingen heeft een spreiding met $N - 1$ andere metingen. (Eigenlijk tel je nu dubbel, want het verschil (de spreiding) tussen x_i en x_j is natuurlijk hetzelfde als het verschil tussen x_j en x_i . Dat maakt voor de berekening niet uit, want als we straks gaan middelen, delen we ook door twee keer zo veel.)

Zoals de spreiding van twee (2) getallen maar één ($2 - 1$) waarde heeft, heeft de spreiding van N getallen maar $N - 1$ waarden.

24. Beste schatting (Pythonsimulatie)

Laten we eens uitgaan van een situatie waar niets misgaat en alles normaal is. Normaal verdeeld zelfs: dat de metingen gemodelleerd kunnen worden als een gausscurve met een gemiddelde van 5 V en een standaardafwijking van 1 V. Als je dan heel veel metingen doet, krijg je een steekproefgemiddelde dat heus wel dicht bij de 5 V zal liggen. Maar ja, datzelfde geldt voor de mediaan. Een normale verdeling is nu eenmaal symmetrisch om het gemiddelde. Niet alleen het gemiddelde is gelijk aan μ , de mediaan is dat ook.

Met en Weten is niet in eerste instantie bedoeld voor studenten die steengoed willen worden in statistiek. Het mikt meer op mensen die beter willen beseffen waar ze mee bezig zijn. Een goed onderbouwd gevoel is minstens zo handig als een sluitend wiskundig bewijs.

Wat zou je kunnen doen om de stelling van Barford dat je *in the absence of arguments to the contrary* voor het gemiddelde moet kiezen? Kritisch nadenken vraagt om tegenargumenten. Soms helpt het om dan een Pythonsimulatie te programmeren. Die taal leent zich daar prima voor, omdat je eenvoudig een reeks normaal verdeelde toevallige getallen kunt genereren. Met de module `statistics` kun je vervolgens allerlei analyses doen.

Op de volgende pagina staat een programma dat een aantal keer negen waarden genereert uit een standaardnormale verdeling. Het kijkt vervolgens telkens hoe ver het gemiddelde en de mediaan van nul af liggen. (Het “echte” gemiddelde van de standaardnormale verdeling is immers gelijk aan nul.)

Als je het programma uitvoert, vind je twee getallen: 0,30 en 0,25. Maar als je het nóg een keer doet, vind je twee andere getallen: 0,33 en 0,27. Doe je het een derde keer, dan kom je op 0,32 en 0,26. Bewijst dit dat het gemiddelde altijd dichterbij de “werkelijke waarde” zit dan de mediaan? Nee. Echt bewijzen doe je met een simulatie sowieso niets. Maar als je in de code de variabele `aantal_herhalingen` flink hoger maakt, zul je zien dat er op een gegeven moment wél steeds bijna hetzelfde uitkomt. Bij tienduizend herhalingen wordt de spreiding kleiner. Bij honderdduizend herhalingen varieert de tweede decimaal nog een beetje. Bij een miljoen herhalingen gaat de simulatie wel erg lang duren. En bij tien miljoen heb je echt geen zin meer om erop te wachten. Het goede nieuws is: dan komt er wel telkens in twee decimalen hetzelfde uit, namelijk 0,32 en 0,27.

Was het maar zo’n feest! Ik deed deze simulatie tien miljoen keer en ging koffiezetten, kwam terug en keek naar de resultaten in zo veel mogelijk decimalen. Mijn programma had 0,3250020951149268 en

0,2660488796561978 gevonden. Leuk, maar dat het derde cijfer van het eerste getal een 5 is, maakt natuurlijk wel dat je heus niet zomaar precies hetzelfde vindt in twee decimalen. En dan nog iets: blijft mijn programma wel goed rekenen als ik miljoenen optellingen doe? Simulaties helpen best wel om iets in te zien. Bijvoorbeeld dat ook simuleren zijn grenzen kent.

O ja, natuurlijk deed ik het nog een keer. Wéér koffiezetten en andere waarden vinden: 0,32483231377518385 en 0,265976253045457. Maar waar was het ook alweer allemaal om te doen? Dat de mediaan gemiddeld meer afwijkt van de werkelijke waarde dan het gemiddelde.

Hebben we nu “bewezen” dat het gemiddelde een betere maat is dan de mediaan. Nee. Voor normaal verdeelde processen lijkt dat wel het geval te zijn. Maar we hebben nu gekeken naar de *absolute waarden* van de verschillen. Zou het niet beter zijn om de *kwadraten* van de verschillen te vergelijken? Of, erger nog: de mediaan van de verschillen.

Zucht.

Zo kun je bezig blijven. Maar met ons Pythonprogramma is het niet veel werk om ook dat te simuleren.

```
import numpy as np
import statistics as st

mediaan_som = 0
gemiddelde_som = 0
aantal_waarden = 9
aantal_herhalingen = 1000

for i in range(aantal_herhalingen):
    waarden = np.random.normal(size=aantal_waarden)
    mediaan = st.median(waarden)
    gemiddelde = st.mean(waarden)

    mediaan_som = mediaan_som + abs(mediaan)
    gemiddelde_som = gemiddelde_som + abs(gemiddelde)

print(round(mediaan_som / aantal_herhalingen, 2))
print(round(gemiddelde_som / aantal_herhalingen, 2))
```

25. Fout gebruik van de term ‘error’

Misschien is *error* (fout) wel de meest verwarrende term die je tegenkomt in het vakgebied van metingen en onzekerheden, van *Metten & Weten* dus. Een meting heeft (vrijwel) altijd een fout, ook als je niets verkeerd hebt gedaan.

Herman Berendsen (1997) gaf zijn oorspronkelijk in het Nederlands uitgegeven boek de titel: *goed meten met fouten*. Je ontkomt niet aan de fout, zelfs niet als je alles goed doet. Maar je kunt het natuurlijk óók verkeerd doen.

In dit vakgebied is een *error* (fout) **NIET** een vergissing of een misser.

Kirkup & Frenkel (2006) zetten in hoofdstuk 3 van hun boek wat begrippen op een rij. In de tabel hieronder staan de Nederlandse vertalingen van die begrippen. De onderste is toegevoegd en staat niet in de genoemde tabel.

measurement	meting
measurand	meetgrootheid
value	waarde
true value	werkelijke waarde
best estimate	beste schatting
error	fout
random error	toevallige fout
systematic error	systematische fout
accuracy	nauwkeurigheid
precision	precisie
uncertainty	onzekerheid
accepted value	geaccepteerde waarde

Een paar van die begrippen halen we nu naar voren.

De waarde (*value*) bestaat uit een getalswaarde én (meestal) een eenheid. Het getal zelf is nog niet de waarde. (Sterker nog: zonder eenheid is de uitkomst in feite waardeloos én zoals eerder gezegd: FOUT.)

Als iemand op een feestje aan je vraagt hoe oud je bent, mag je heus wel ‘twintig’ antwoorden, als dat waar is. Zolang je in de context van dit vakgebied maar netjes ‘20 jaar’ zegt.

De werkelijke waarde (*true value*) is datgene dat we willen weten en meten, en dus niet kennen, zelfs niet na de meting. Daarmee samenhangend kan iets gezegd worden over dat verwarrende begrip *error*.

Dit is de definitie die Kirkup & Frenkel geven:

$$\text{error} = \text{measured value} - \text{true value}$$

Dan kennen we de fout dus niet! We kennen wel de gemeten waarde, maar niet de werkelijke waarde. En dus ook het verschil niet.

Stel dat iemand inderdaad 1,83 meter lang is. En dan ook nog eens *exact* 1,83 meter. Het is moeilijk om je daar wat bij voor te stellen, omdat je bij exact moet denken aan 1,830000000000000000 meter, maar dan met nog veel meer nullen. Oneindig veel dus, niet die slordige zestien die hier staan. Ook al zitten we op een miljardste van een molecuul: het is nog niet exact.

Goed, voor het gedachtenexperiment: iemand is exact 1,83 meter en ik doe een meting. Hoe groot schat ik dan de fout, de *error*? Je zou kunnen zeggen: 2 cm. Maar is dat niet je *onzekerheid*?

Antwoord op deze retorische vraag: ‘Ja, dat is je onzekerheid!’

De fout (*error*) schat ik namelijk op exact nul. Niet omdat ik denk dat ik geen fout maak, maar omdat het nu eenmaal mijn beste schatting is. Als ik dacht dat het een centimeter meer was, had ik wel meer gezegd.

Zinniger is het om een schatting te maken van de *absolute fout*, het *absolute verschil* tussen de gemeten waarde en de werkelijke waarde. Het zou vreemd zijn als ik die op 0,0 cm zou schatten. Het is heus wel meer dan dat.

Taylor (1997) schrijft in de eerste paragraaf van hoofdstuk 1 het volgende.

1.1 Errors as Uncertainties

In science, the word error does not carry the usual connotations of the terms mistake or blunder. Error in a scientific measurement means the inevitable uncertainty that attends all measurements. As such, errors are not mistakes; you cannot eliminate them by being very careful. The best you can hope to do is to ensure that errors are as small as reasonably possible and to have a reliable estimate of how large they are. Most textbooks introduce additional definitions of error, and these are discussed later. For now, error is used exclusively in the sense of uncertainty, and the two words are used interchangeably.

Tsja. *The two words are used interchangeably*. Hij geeft een andere definitie dan Kirkup & Frenkel. En het is ook nogal andere taal dan de GUM gebruikt.

In this Guide, great care is taken to distinguish between the terms “error” and “uncertainty”. They are not synonyms, but represent completely different concepts; they should not be confused with one another or misused.

Kunnen ze niet allebei gelijk hebben? Het duo Kirkup & Frenkel enerzijds en de GUM anderzijds? In een poging om ze op één lijn te krijgen, zou je kunnen zeggen dat je met ‘fout’ één concreet geval van een bepaalde meting kunt bedoelen (GUM) of een kansvariabele met mogelijke uitkomsten. Die kansvariabele druk je dan uit in alleen de standaardafwijking, om het gemiddelde toch nul is. (Vergeef me deze verzoeningspoging, het wordt er niet eenvoudiger op. In *Metten & Weten* volgen we de definitie van de GUM.)

De meeste termen in bovenstaande tabel gaan we nog wel tegenkomen. Over *measurand* (gemeten grootheid) moet nog wel iets worden opgemerkt. Als je een snelheid bepaalt door een afstand en een tijd te meten en daarna een deling uit te voeren, dan noem je die snelheid nog steeds de *measurand*.

Oefening

1. Geef van onderstaande begrippen aan of ze een onzekerheid hebben.

measurement	meting
true value	werkelijke waarde
best estimate	beste schatting
accuracy	nauwkeurigheid
precision	precisie
uncertainty	onzekerheid
accepted value	geaccepteerde waarde

Notatie

26. Afronden

Er zijn verschillende manieren om een getal af te ronden. Welke je neemt, hangt af van het doel dat je hebt. We noemen het laatste cijfer van het afgeronde getal het *relevante cijfer*.

Een gebruikelijke manier van afronden kijkt naar het eerste cijfer na het relevante cijfer. Als dat een 5 of hoger is (een van de cijfers 5 t/m 9), dan wordt het relevante cijfer met 1 verhoogd. Als het lager is dan 5 (een van de cijfers 0 t/m 4), dan blijft het relevante cijfer onveranderd. Zo worden cijfers op schoolrapporten vastgesteld. Een 7,4 wordt een 7. Een 7,5 wordt een 8. Bij lage cijfers is die afronding relatief groot. Wie een 1,4 haalde krijgt een 1. Dat is 29% lager dan wat hij of zij verdient. Wie een 9,4 ziet afgerond worden op een 9 krijgt maar 4,3% minder. Toch ‘voelt’ juist dat verschil heel groot. Als je één tiende meer had gehad, was het een 10 op je eindlijst geworden.

Eigenlijk is het vreemd om dit soort voorbeelden met percentages te berekenen. Als je minder dan een 2 hebt, wat doet het cijfer er dan nog toe? Maar aan de ‘bovenkant’ van de scores is de score om de afrondingsmarges des te groter. Wie een 9,4 had, zat maar 6% van een 10 af. Die afronding maakt dat het 10% wordt: meer dan anderhalf keer zo veel.

Soms rond je alles af *naar boven*. Voor een excursie met 204 studenten wil je bussen huren en er kunnen 60 personen in één voertuig. Dan zou je er 3,4 nodig hebben. Als je dat zou afronden naar beneden, kunnen er 24 studenten niet mee. Daarom moet je afronden naar boven.

Als je projectgroepen wilt maken die uit minstens vijf studenten bestaan – maar liever niet veel meer dan dat – deel je het aantal studenten in de klas door 5 en rond je af *naar beneden* als je wilt uitrekenen hoeveel groepen er komen. Een klas van 32 studenten wordt dan in 6 groepen verdeeld. Je hebt dan vier groepen van vijf studenten en twee groepen van zes.

De ‘gewone’ manier van afronden heeft een nadeel. De 5 zit namelijk precies tussenin. Als je het getal 7,5 naar 7 afrondt, is het verschil met de oorspronkelijke waarde net zo groot als wanneer je naar 8 afrondt. Voor een 7,6 geldt dat niet: die zit 0,4 punt van de 8 vandaan en meer (0,6) van de 7. Ook voor 7,5000001 geldt dat het verschil met 8 kleiner is dan met 7. Maar met exact 7,5 is het even groot.

Als je uitgaat van een getal met één decimaal dat je wilt afronden op een geheel getal, dan zijn er tien mogelijkheden. Die ene decimaal waar je van af wilt, kan namelijk 0 t/m 9 zijn.

Voor een 0 geldt dat er niets af te ronden valt. Bij een 1 als decimaal raak je 0,1 punt ‘kwijt’, maar dat wordt gecompenseerd door de 9 als decimaal, want daarbij ‘win’ je juist 0,1 punt. Zo compenseert de 2 ook de 8 en de 3 de 7, net als de 4 tegenover de 6. Alleen voor de 5 geldt dat – als je die altijd omhoog zou afronden – nooit gecompenseerde wordt.

Bij *banker's rounding* (ook wel *Dutch rounding* of *round-to-even* genoemd) wordt het cijfer 5 op een aparte manier gebruikt. Je rondt namelijk niet af naar boven, maar naar het dichtstbijzijnde even cijfer. Een 5,5 rond je dus af op een 6, maar een 4,5 op een 4. Het voordeel van deze aanpak is dat de verwachtingswaarde van het gemiddelde van afgeronde getallen (uniform verdeeld) gelijk zal zijn aan de verwachtingswaarde van het gemiddelde van niet-afgeronde getallen.

Hoe groot is het verschil ten opzichte van ‘gewoon’ afronden? Om dat te berekenen, moet je inzien dat er twintig verschillende mogelijkheden zijn: alle tien cijfers 0 t/m 9 kunnen als eerste cijfer achter de komma staan. Het getal vóór de komma heeft ook tien mogelijkheden, maar het enige wat van belang is, is of er een even of een oneven getal staat. Dat maakt dus dat er twee keer tien mogelijkheden zijn. Bij negentien van die twintig mogelijkheden is er geen verschil tussen ‘gewoon’ afronden en *banker's rounding*; alleen een 5 met een even cijfer ervoor maakt verschil. Die wordt namelijk *omlaag* afgerond in plaats van omhoog. Het verschil in de afgeronde getallen is dan 1: de 4,5 wordt 4 in plaats van 5. Omdat dit in één op de twintig gevallen voorkomt, zal het gemiddelde bij *banker's rounding* $1/20 = 0,05$ punt lager liggen. Eigenlijk moet je dat natuurlijk omkeren: het gemiddelde van alle mogelijke getallen is bij *banker's rounding* namelijk precies gelijk voor onafgeronde en afgeronde getallen.

Er wordt in boeken verschillend omgegaan met afrondingen. Grabe (2014) rondt het cijfer 5 altijd af naar boven. Kirkup & Frenkel (2006) gebruiken *banker's rounding*, evenals Bevington & Robinson (2002).

Grabe kiest ervoor om onzekerheden *altijd* naar boven af te ronden. Niet alleen de 5 dus, maar *alle* cijfers. Een voorbeeld in zijn boek is het getal (met onzekerheid) $119748,8 \pm 123,7$. Hij noteert dat als 119750 ± 130 . In *Metten & Weten* zouden we kiezen voor $(1,1975 \pm 0,0012) \cdot 10^5$. Dat ziet er iets ingewikkelder uit, maar volgt wel de striktere regels rondom significantie. Bovendien is er een verschil in afronding: wij maken 12 van 12,37. En we zijn niet de enigen: Taylor (1997) doet het ook zo, net als Hughes & Hase (2010).

Je kunt niet zeggen dat sommigen het fout doen en anderen goed, maar het is wel belangrijk om te snappen wat er achter de verschillende keuzes zit. Strikt genomen is *banker's rounding* het meest zuiver. Aan de

andere kant: waarom moeilijk doen over $(1,45 \pm 0,32)$ A dat eigenlijk zou moeten worden afgerond op $(1,4 \pm 0,3)$ A en niet op $(1,5 \pm 0,3)$ A, terwijl de onzekerheid in het cijfer achter de komma al zo groot is?

Sommige boeken kiezen zelfs voor $(1,5 \pm 0,4)$ A, omdat je de onzekerheid ‘voor de zekerheid’ naar boven afrondt. Die keuze gaat ervan uit dat de onzekerheid een soort absolute grens aangeeft. Maar als je met onzekerheid bedoelt dat er een redelijk grote *kans* is dat de werkelijke waarde in dat bereik ligt, dan is het zinvol om zo weinig mogelijk af te ronden.

Als je te gedetailleerd gaat kijken, verlies je het gevoel voor verhoudingen. Een onzekerheid van 320 mA op een stroomsterkte van 1,45 A is een relatieve onzekerheid van 22%. Dat is een ander verhaal dan 1% of 0,01%. Je kunt er bijna een kwart naast zitten.

Oefeningen

1. Bereken of beredeneer wat het effect zou zijn als *banker's rounding* zou worden toegepast op tentamencijfers. Wat voor verschil zou dat maken op de gemiddelde cijfers van studenten?

Wat staan er toch veel lollige dingen op *Wikipedia*! De pagina over het getal pi vermeldt deze zin:

Wat u door 'n domme ezelsbrug te kennen, immer met gemak onthoudt.

Als je het aantal letters van elk van deze twaalf woorden opschrijft en een komma noteert na het eerste cijfer, krijg je pi in elf decimalen: 3,14159265358. Als je het zou afronden, moet die laatste een 9 zijn.

Engelstaligen schijnen onderstaande zin te gebruiken als *donkey bridge*. (Zo heet een Engelse ezelsbrug niet en alcohol is slecht voor je.)

How I need a drink, alcoholic of course, after the heavy chapters involving quantum mechanics.

2. Je kunt het getal 3,14159265358 afronden op 1 of 2 decimalen, maar ook op 3, 4, 5, 6, 7, 8, 9 of 10. Voor welke van die tien opties geldt dat het uitmaakt of je *banker's rounding* of gewoon afronden gebruikt?
3. Rond het hexadecimale kommagetal 1A,F84 af op één cijfer achter de komma. Doe dat via “gewoon” afronden én via *banker's rounding*.
4. Rond het binaire kommagetal 101,010 af op één cijfer achter de komma. Doe dat via “gewoon” afronden én via *banker's rounding*.

27. Significante cijfers

Om metingen met hun onzekerheden goed te kunnen noteren, moeten we weten wat *significante cijfers* zijn. In Vlaanderen noemen ze dat *beduidende* cijfers. Het zijn cijfers die van betekenis zijn. Dat geldt voor alle cijfers behalve *voorafgaande nullen*. Voorafgaande nullen zijn alle nullen die (vanaf links gezien) vóór een ander cijfer staan.

Het getal 0 heeft één significant cijfer. Het staat niet vóór een ander cijfer.

Voor nullen ‘aan het eind’ geldt dat ze altijd significant zijn als ze na de komma staan. Als je geen komma hebt, is er eigenlijk geen duidelijkheid over nullen aan het eind. Je kunt van 1000 zeggen dat het vier significante cijfers heeft, maar één is ook verdedigbaar. Van 1,000 geldt dat het sowieso vier significante cijfers heeft. Anders hadden we wel wat nullen weggelaten.

Tabel 27.1 Voorbeelden van getallen met significante cijfers

getal	niet-nul	significante nullen	significante cijfers	getal grijs en onderstreept
2023	3	1	4	<u>2023</u>
3,14	3	0	3	<u>3,14</u>
0,98	2	0	2	<u>0,98</u>
1,0001	2	3	5	<u>1,0001</u>
0,0060	1	1	2	<u>0,0060</u>
7,50	2	1	3	<u>7,50</u>
0,000001	1	0	1	<u>0,000001</u>
$7,20 \cdot 10^6$	2	1	3	<u>$7,20 \cdot 10^6$</u>
0	0	1	1	<u>0</u>
0,00	0	1	1	<u>0,00</u>
1000	1	0, 1, 2 of 3	1, 2, 3 of 4	1000
$1,00 \cdot 10^3$	1	2	3	<u>$1,00 \cdot 10^3$</u>

De laatste twee regels in deze tabel geven samen een voorbeeld van het geval dat niet duidelijk is. Als iemand het getal ‘1000’ opschrijft, weet je zonder toevoeging niet welke cijfers significant zijn. Wellicht heeft iemand 1,0 kΩ en zijn het tweede en derde cijfer achter de komma helemaal niet bekend. Als je dan liever in ohms rekent, moet je van die 1,0 kΩ toch 1000 Ω maken. Het onderste voorbeeld is de oplossing voor dit probleem als je geen gebruik wilt maken van voorvoegsels als kilo, mega etc.

Er staat in deze tabel ook twee keer ‘nul’. Niks dus, noppes. Eén keer in de vorm ‘0’ en één keer in de vorm ‘0,00’. Hoeveel significante cijfers zijn dat?

Stel dat het geld betreft, of een spanning...

Als de Belastingdienst zegt dat je € 0 terugkrijgt, kun je vermoeden dat er centen in je nadeel zijn afgerond. (Wat niet waar is, want er wordt altijd – volgens de regels althans – in je voordeel afgerond.) Als je leest dat je € 0,00 terugkrijgt, gaat het hooguit om de afronding van tienden van centen.

Als je een spanning meet van 0 V, dan kunnen er nog tienden van een volt (positief of negatief) aanwezig zijn. Die zie je niet op de meter. Maar wat als je 0,000 V meet? Of 0 mV? Dan weet je dat er hooguit een spanning van een paar tienduizendste volt aanwezig is.

Toch is het aantal significante cijfers bij al deze spanningen gelijk aan 1. De metingen verschillen wel in onzekerheid, maar niet in significante cijfers.

Oefeningen

Noem het aantal significante cijfers in onderstaande waarden of grootheden.

1. 0 A
2. 1
3. 0,1 A
4. 0,10 mA
5. 0,0000000000000000000010
6. € 1,11
7. € 0,00
8. 0,01100
9. 1000000
10. $1,0 \cdot 10^6$
11. 1.000.000,00

28. Significante cijfers in berekeningen

Wat moet je doen met significante cijfers in berekeningen? In tussentijdse resultaten laat je in principe alle cijfers staan. In ieder geval neem je er zoveel mogelijk mee in je berekening. Maar wat doe je met het eindresultaat?

Stel je wilt een snelheid meten. Je meet de afstand waarover een voorwerp zich verplaatst en komt op 2,235 m. Dat is dus een afstand gemeten in *millimeters*. Een extra cijfer aflezen lukt niet, maar op millimeters klopt het wel. Het aantal significante cijfers is dus vier. Je meet de tijd die het voorwerp nodig heeft voor de verplaatsing en komt op 2,5 s. De tijdmeting die je toepast, staat niet toe om nog een cijfer achter de komma te zetten. Je hebt dus twee significante cijfers. Nu kun je de snelheid uitrekenen.

$$v = \frac{d}{t} = \frac{2,235 \text{ m}}{2,5 \text{ s}} = 0,894 \text{ m/s}$$

In dit geval levert een deling van een grootheid met vier significante cijfers door een grootheid met twee significante cijfers een resultaat op met drie decimalen. Maar als je 2,4 s had gemeten, zou er 0,93125 m/s uitgekomen zijn. En als je in milliseconden had kunnen meten en je was op 2,235 s uitgekomen, dan was je antwoord 1 m/s geweest: één cijfer dus.

Los van het feit dat het aantal cijfers in het antwoord niet telkens overeenkomt: het is een wat vreemde gedachte dat je door te delen door iets met weinig significante cijfers er net zo veel of meer overhoudt. Als je niet ‘beter’ dan in hele seconden had kunnen meten, zou je door 2 s hebben gedeeld. Dan was de uitkomst 1,1175 m/s geweest. Maar liefst vijf significante cijfers.

Hoe goed ken je de waarden waarmee je rekt eigenlijk? Ervan uitgaande dat alle significante cijfers een correcte waarde hebben, is de echte afstand minstens 2,2345 m en hoogstens 2,23549999 m. Anders zou je niet op deze cijfers hebben afgerond. Voor de gemeten tijd van 2,5 s geldt dat het minimaal 2,45 s was en maximaal 2,549999 s. (Na die vier negens komen er nog oneindig veel.)

Je zou dus kunnen zeggen dat de relatieve onzekerheid in de afstand gelijk is aan:

$$\frac{\delta d}{d} = \frac{0,0005 \text{ m}}{2,235 \text{ m}} = 0,02237\%$$

De relatieve onzekerheid in de tijd is gelijk aan:

$$\frac{\delta t}{t} = \frac{0,05 \text{ s}}{2,5 \text{ s}} = 2\%$$

Kirkup & Frenkel (2006) geven als vuistregel voor het aantal significante cijfers dat je moet kiezen voor een aantal dat zorgt dat de relatieve onzekerheid zo dicht mogelijk in de buurt van de grootste van die twee relatieve onzekerheden komt te liggen. Als je opties in een tabel weergeeft, kom je dan op deze waarden.

Waarde	Onzekerheid	relatief
0,894 m/s	0,0005 m /s	0,056%
0,89 m/s	0,005 m/s	0.56%
0,9 m/s	0,05 m/s	5,6%

De auteurs geven dezelfde aanpak voor optellen en aftrekken. Ze gebruiken als rekenvoorbeeld het optellen van de massa van een koperen vat met water (1,5778 kg) en de massa van alleen het vat (0,582 kg). Ze komen dan op percentages van 0,003% en 0,05%.

Wat nu als je een olifant hebt gewogen met een muis op zijn rug? Je meet dat ze samen 4162 kg wegen. De muis kun je eraf halen en daarna wegen op een andere weegschaal. Je stelt vast dat hij een massa heeft van 19 g. De relatieve onzekerheden in de metingen zijn $0,5 / 4162 \approx 0,012 \%$ en $0,5 / 19 = 2,6\%$. Dat zou je dus volgens Kirkup & Frenkel van die 2,6% moeten uitgaan...

De olifant zonder de muis weegt $4162 - 0,019 = 4161,981$ kg. Met een onzekerheid van 2,6% zou de onzekerheid in de massa van de olifant alleen dus 110 kg zijn. Hier lijkt iets niet te kloppen: je gezonde verstand zegt dat de massa van de muis te verwaarlozen is.

Een veel eenvoudigere vuistregel is:

- Gebruik bij optellen en aftrekken het kleinste aantal *cijfers achter de komma*.
- Gebruik bij vermenigvuldigen en delen het kleinste aantal *significante cijfers*.

Bij het voorbeeld van de olifant en de muis had de eerste nul cijfers achter de komma en de tweede drie. Die 4161,981 kg rond je dus af op 4162 kg. Wat we al dachten: de massa van de muis is verwaarloosbaar.

Bij het voorbeeld van de snelheid hadden we twee significante cijfers in de tijd en vier in de afstand. Volgens de simpele vuistregel kom je dan op twee significante cijfers. Die 0,894 m/s moet dus worden genoteerd als 0,89 m/s.

Wat nu als we in milliseconden hadden gemeten en die 2,235 s hadden gevonden? Zou die 1 m/s dan nog steeds 1 m/s zijn? Ja, maar wel met meer decimalen. Als tijd en afstand beide in vier significante cijfers waren opgegeven, zou het dus 1,000 m/s moeten zijn. Let op: dat zijn *drie* decimalen en *vier* significante cijfers.

Oefeningen

Bereken onderstaande sommetjes en geef de antwoorden in het juiste aantal significante cijfers.

1. $456 + 3,5$
2. $100 - 99,9$
3. $3,1415 \times 100$
4. $12,0 / 2,54$
5. $7,21 \times 8,66 \times 2$
6. $1,2345^2$
7. de lichtsnelheid in kilometer per uur

29. Significante cijfers in onzekerheden

Hoeveel significante cijfers wil je gebruiken voor een onzekerheid? Stel dat je de versnelling van de zwaartekracht hebt gemeten en je komt uit op

$$g = (9,81643552 \pm 0,03244913) \text{ m/s}^2$$

Hoeveel zin heeft het dan om acht cijfers achter de komma te geven? De waarde ligt in Nederland tussen de 9,78 en 9,84 m/s². De cijfers die erna komen, hebben dus geen praktische betekenis.

Je zou ook kunnen uitrekenen wat de relatieve fout in g is. Als je bovenstaande twee getallen op elkaar deelt, kom je op 0,33%. Neem je de onzekerheid in één significant cijfer, dan kom je 0,3%. Als je naar die twee getallen kijkt (0,33% en 0,30%) dan lijkt dat een relevant verschil: die 0,03% (procentpunt) is namelijk 10% (procent) van 0,30%. Maar heeft het nu zin om te zeggen dat de waarde ligt tussen 9,784 en 9,848 m/s², in plaats van tussen 9,78 en 9,84 m/s²? Of de absolute onzekerheid nu 0,064 of 0,06 m/s², maakt ook niet echt uit. Het gaat om indicatie, een indruk van hoe groot het gebied is waarin de waarde waarschijnlijk ligt.

Er is nog een reden om niet met twee significante cijfers te werken: vaak weet je de onzekerheid zelf helemaal zo zeker niet. We komen daarop terug in hoofdstuk 77, als we daar statistische berekeningen bij halen.

Als regel geven we onzekerheden op met één significant cijfer. Daar is een uitzondering op: als dat ene cijfer een 1 is, voegen we er eentje toe om op twee significante cijfers te komen. Dit wordt zo aangeraden door Hughes & Hase (2010) en door Taylor (1997). Fornasini (2008) zegt dat het er nooit meer dan twee moeten zijn en dat één soms al voldoende is. Volgens Kirkup & Frenkel (2006) is het ongebruikelijk om meer dan twee significante cijfers op te geven bij een onzekerheid. Grabe (2014) kiest ervoor om zowel bij een 1 als bij een 2 een extra cijfer op te geven.

Is hier dan geen norm voor? Nee. Van de vijf boeken die we hierboven aanhalen, doen twee ervan precies hetzelfde. (Toevallig precies de twee die bij de opleiding Mechatronica aan de Haagse Hogeschool op de boekenlijst hebben gestaan...) De andere drie boeken zijn daar niet echt strijdig mee. Nou ja, die van Grabe wel. Volgens het boek van Grabe hoort er een tweede cijfer achter zowel de 1 als de 2. De gedachte erachter is hetzelfde, maar het is anders.

Dat je bij een 1 een extra cijfer wilt opgeven, heeft te maken met het feit dat afrondingen bij dat cijfer relatief nogal fors kunnen zijn. Stel je voor dat je

een lengtemeting hebt met een onzekerheid van 0,0149952 m... Als je dat noteert als 0,01 m, dan heb je in feite 15 mm afgerond naar 10 mm. De berekende onzekerheid is dan maar liefst 50% meer dan wat je opgeeft. Dat leed zou te verzachten zijn door naar boven af te ronden op 2, maar dan is het nog steeds 25%. Bij hogere cijfers is dat relatieve verschil kleiner. Overigens kan een 2 ook een afronding zijn van 2,499 en dan zit je ook met die 25%. Toch volgen we in dit boek de lijn van Taylor en van Hughes & Hase.

Natuurconstanten worden in de literatuur vrijwel altijd met twee significante cijfers in hun onzekerheden opgegeven, ook als de eerste geen 1 is. Nu zijn de metingen daarvan vaak zo extreem goed, dat twee cijfers in de onzekerheid ook niet heel vreemd zijn.

Oefeningen

Noteer de volgende onzekerheden met het juiste aantal significante cijfers.

1. 5,94 s
2. 14,56 m
3. $1,55 \cdot 10^6$ kg
4. 2,51 A
5. 1,91 K
6. 0,1999 V

30. Noteren van onzekerheden

Je hebt een meting pas goed genoteerd als de volgende punten correct zijn.

1. de eenheid van de beste schatting en de onzekerheid
2. de waarde van de onzekerheid in het juiste aantal cijfers
3. de waarde van de beste schatting in het juiste aantal cijfers
4. eventuele machten van 10 en voorvoegsels zoals ‘milli’
5. een punt of een komma op de juiste plaats
6. eventuele spaties om de waarde duidelijker te maken
7. eventuele haakjes om het geheel duidelijker te maken

Bewust is in dit lijstje de eenheid op nummer 1 gezet. Zonder een eenheid heb je namelijk niets, ook kun je het in praktijk vaak zelf wel invullen. Als iemand zegt dat hij of zij een houten ondergrond van “twee bij twee” gaat aanleggen onder een meetopstelling, dan zullen dat heus wel meters zijn. Als dezelfde persoon zegt: “Ik ben 34”, dan is het vast een leeftijd in jaren. “Ik ben één zesentachtig” is meestal een lengte in meters (1) en centimeters (86).

Met stip op 2: de waarde van de onzekerheid. Ook daarvoor geldt: als je die niet kent, weet je ook weinig. “Ik heb een voorlopig cijfer voor je tentamen: je hebt een 6,2.” “Joepie!” “En de onzekerheid in het cijfer is 3,5 punt.”

Er is nog een reden om met de onzekerheid te beginnen. Het juiste aantal cijfers daarin is namelijk bepalend voor de notatie van de beste schatting. Dat is ook wel logisch: de onzekerheid zegt namelijk iets over de waarde ervoor, namelijk hoe goed je die kent. Het omgekeerde is niet het geval.

Stel dat je een multimeter hebt met te veel cijfers achter de komma. Je meet daarmee een weerstand en komt uit op de volgende waarde.

$$R = (122,1437668 \pm 0,0034220575623) \Omega$$

Hoe ga je dit in de juiste vorm schrijven? Begin met de eenheden (die kloppen al) en de onzekerheid. Die noteer je met één significant cijfer.

$$\delta R = 0,003 \Omega$$

Let op: de onzekerheid in de gemeten weerstand is hier dus in *één significant cijfer* en in *drie decimalen*. Om de beste schatting daarmee te laten kloppen, moet die waarde ook in drie decimalen worden genoteerd. We krijgen dus:

$$R = (122,144 \pm 0,003) \Omega$$

Het derde cijfer achter de komma wordt een 4, omdat er na de 3 een 7 stond.

Is de weerstand zelf nu ook in één significant cijfer gegeven? Nee, de weerstand zelf heeft zes significante cijfers. Als de onzekerheid tien keer zo klein was geweest (0,0003 Ω dus), waren dat zelfs zeven cijfers geweest. De beste schatting had dan namelijk als 122,1438 Ω moeten worden genoteerd. En was de onzekerheid honderd keer zo groot (0,3 Ω dus), dan bleven er vier significante cijfers over. Drie voor de komma en één erachter.

Stel nu dat het eerste significante cijfer van de onzekerheid in de weerstand niet een 3 maar een 1 was geweest, en de andere cijfers hetzelfde; wat zou dan het resultaat zijn geworden? De onzekerheid zou in dat geval in één decimaal extra worden gegeven (twee significante cijfers) en de beste schatting ook. De afronding pakt dan wel anders uit. In dat geval komen we op:

$$R = (122,1438 \pm 0,0014) \Omega$$

Dit zijn twee voorbeelden met nogal kleine onzekerheden. Stel nu eens dat de onzekerheid veel groter was.

$$\delta R = 10,00 \Omega$$

We moeten dan de onzekerheid in twee significante cijfers schrijven, omdat het eerste cijfer in δR een 1 is. Dan blijft een waarde over zonder cijfers achter de komma. Voor de beste schatting geldt dan hetzelfde.

$$R = (122 \pm 10) \Omega$$

Maar wat nu als de onzekerheid nog drie keer zo groot was? Het is verleidelijk om dan in plaats van 10 gewoon 30 te schrijven, maar dat mag niet: een 3 is geen 1, en moet dus met één significant cijfer worden genoteerd. Een oplossing die je in sommige boeken wel ziet, is dat de laatste 2 van 122 wordt weggelaten. Of beter gezegd: omgezet in een 0.

$$R = (120 \pm 30) \Omega$$

Het nadeel hiervan is dat je volgens de regels van significantie nog steeds te veel cijfers hebt. We kunnen dit oplossen door met machten van 10 te werken, of te kiezen voor een voorvoegsel in de eenheid.

$$R = (0,12 \pm 0,03) \text{ k}\Omega$$

Wie niet van onnodige komma's houdt, wordt hier niet blij van. Toch heeft deze notatie een voordeel: de $\text{k}\Omega$ maakt duidelijk dat je het zo heel zeker niet weet.

We gaan nu een spanning meten met de multimeter. Dit is het meetresultaat.

$$V = 3,7502 \text{ V} \pm 42 \text{ mV}$$

De eenheden kloppen nu dimensioneel wel met elkaar, maar het voorvoegsel *m* (milli, een duizendste) maakt dat de getallen niet met elkaar overeenkomen. Er moet dan gekozen worden tussen volts en millivolts. Vanwege de grootte van de beste schatting liggen volts hier het meest voor de hand. Om fouten te voorkomen, voeren we die eerste stap apart uit.

$$V = 3,7502 \text{ V} \pm 0,042 \text{ V}$$

Het lijkt een simpele zet, maar je gaat met het verschuiven van de komma heel snel de mist in. Van millivolts naar volts is een factor 1000. Je moet de komma dus drie plaatsen naar links zetten. Daarna is het nog een kwestie van het juiste aantal significante cijfers kiezen en de haakjes om de getallen.

$$V = (3,75 \pm 0,04) \text{ V}$$

Naast voorvoegsels kunnen er ook machten van 10 zijn die maken dat de beste schatting en de onzekerheid nog moeten worden ‘uitgelijnd’. Dat zie je in onderstaand voorbeeld.

$$l = 0,01773452 \cdot 10^{-2} \text{ mm} \pm 6,222 \cdot 10^{-3} \mu\text{m}$$

Een raar geval, maar als oefening des te aardiger. Waar begin je mee? De eerste stap zou kunnen zijn om die machten van 10 weg te werken. Dan krijg je veel cijfers achter de komma, maar dat komt daarna wel goed.

$$l = 0,0001773452 \text{ mm} \pm 0,006222 \mu\text{m}$$

Om ze aan elkaar gelijk te maken, moet er worden gekozen tussen mm en μm . Natuurlijk zou je een ander voorvoegsel kunnen kiezen, maar als je dat doet, maak je twee stappen tegelijk en dat vergroot de kans op fouten. Gegeven de beide getallen liggen micrometers het meest voor de hand. De beste schatting moet met een factor 1000 worden vermenigvuldigd, de komma gaat drie plaatsen naar rechts.

$$l = 0,1773452 \mu\text{m} \pm 0,006222 \mu\text{m}$$

$$l = (177,3452 \pm 6,222) \text{ nm}$$

De laatste stap is het correct noteren van het aantal *significante cijfers* in de *onzekerheid* en het daarmee laten overeenkomen van het aantal *decimalen* in de *beste schatting*. Let in die stap goed op de afronding: in dit geval omlaag, omdat er na de komma een 3 staat.

$$l = (177 \pm 6) \text{ nm}$$

Op toetsen van het vak *Error Budgeting* in de opleiding Mechatronica van de Haagse Hogeschool (locatie Delft) wordt al jarenlang dezelfde vraag 1 gesteld. Die vraag luidt als volgt.

Herschrijf de volgende (berekende) grootheden zodanig dat het aantal significante cijfers klopt met de regels zoals die gegeven zijn in het vak Error Budgeting.

Studenten krijgen vervolgens vijf metingen en onzekerheden die ze op de correcte manier moeten opschrijven. De vraag is één punt waard en per fout gaat er een halve punt af. Als je twee of meer (van de vijf) vragen fout doet, heb je dus nul punten... Streng? Misschien. Maar studenten weten van tevoren dat ze deze vraag krijgen! En de regels zijn gegeven en uitgelegd.

Oefeningen

Herschrijf de volgende (berekende) grootheden zodanig dat het aantal significante cijfers klopt met de regels zoals die gegeven zijn in dit hoofdstuk.

1. lengte = $1,71 \text{ m} \pm 1\frac{1}{2} \text{ cm}$
2. breedte = $200 \text{ mm} \pm \text{een half inch}$
3. oppervlakte = $60,55 \text{ m}^2 \pm 450 \text{ cm}^2$
4. volume = $40,3020 \text{ cl} \pm 10 \text{ cc}$
5. tijd = $59,9876 \text{ s} \pm 5 \text{ ms}$
6. massa = $72,55 \text{ kg} \pm 391 \text{ g}$
7. hoek = $3,1415279 \pm 0,025 \text{ rad}$
8. volume = $98,357 \text{ cm}^3 \pm 246 \text{ mm}^3$
9. frequentie = $50 \text{ Hz} \pm 5\%$
10. temperatuur = $20 \text{ }^\circ\text{C} \pm 2 \text{ }^\circ\text{F}$
11. elektrische stroom = $1,101011001 \text{ A} \pm 1,49 \text{ mA}$
12. volume = $0,44 \cdot 10^{-2} \text{ mm}^3 \pm 7,1 \cdot 10^{12} \text{ nm}^3$
13. spanning = $230,00000500025 \text{ V} \pm 45 \text{ } \mu\text{V}$
14. elektrische stroomdichtheid = $0,04 \text{ A/m}^2 \pm 400 \text{ mA/mm}^2$
15. warmte = $6033 \text{ kJ} \pm 4499 \cdot 10^2 \text{ J}$
16. oppervlakte = $7659 \cdot 10^{-6} \text{ } \mu\text{m}^2 \pm 8,834 \cdot 10^{-1} \text{ nm}^2$
17. kracht = $99,9 \text{ kN} \pm 0,99 \cdot 10^4 \text{ N}$
18. vermogen = $0,22 \text{ kW} \pm 33 \text{ W}$
19. druk = $200 \text{ Pa} \pm 3 \text{ N/dm}^2$
20. weerstand = $2 \cdot 10^4 \text{ } \Omega \pm 0,3 \text{ k}\Omega$

31. De Grote Negen

Iemand heeft een slingerproef gedaan en heeft daarmee geprobeerd om de versnelling van de zwaartekracht op een bepaalde plaats te bepalen.

“Zo dat is aardig gelukt! Ik kom uit op 9.876...”

Au. Hier gaan wat dingen mis.

De berekening resulteerde in 9,87654321 m/s². Wat toevallig! Nee hoor. Het voorbeeld is verzonnen. Natuurlijk zou het erg toevallig zijn als je bij een echte meting op precies alle tien cijfers van het decimale talstelsel uitkomt. Maar in boeken en films gebeuren wel vaker toevallige dingen.

Nee, het punt is *die punt*. In het Nederlands gebruik je een *komma*. (Fout 1)

Maar ook de getalswaarde klopt niet: als je 9,87654321 wilt afronden op drie decimalen, dan kom je uit op 9,877. (Fout 2)

Nu kan een getalswaarde zonder eenheid natuurlijk nooit de versnelling van de zwaartekracht zijn.

“Ja, die liet ik voor het gemak even weg.”

“Wat was je aan het meten dan?”

“De Gravitatieconstante. Dat is hoe sterk de zwaartekracht is. Dus als jij zo nodig eenheden erbij wilt: het was in Newtons.”

Au. Ook hier gaan wat dingen mis. En dan bedoel ik nog niet eens die hoofdletter bij de gravitatieconstante. Dat is een natuurconstante die wel te maken heeft met de zwaartekracht, maar is iets anders dan wat we willen meten: de versnelling van de zwaartekracht op een bepaalde plaats.

De eenheid is niet *newton* (met een kleine letter, dus Fout 3), maar m/s².

Als je die weglaat (Fout 4), zeg je eigenlijk niets. Of iets verkeerds.

De persoon die het experiment uitvoerde is zo eigenwijs om geen komma te gebruiken in een Nederlandse tekst en probeert dat daarom nu op te lossen door een macht van tien te gebruiken.

$$g = 9877 \pm 25 \cdot 10^{-3} \text{ m/s}^2$$

Leuk geprobeerd, maar als je het zo opschrijft, heeft die macht van tien alleen maar betrekking op de *onzekerheid*. Hier moeten toch echt haakjes omheen. (Fout 5) Tenzij je de macht van tien gewoon niet gebruikt. Je hoeft geen haakjes te zetten vanwege die eenheid: die hoort namelijk vanzelfsprekend bij *beide*. Een absolute onzekerheid heeft altijd dezelfde eenheden als de meetwaarde zelf.

$$g = 9,88 \pm 0,25 \text{ m/s}^2$$

En hier gaat wéér iets mis. Je maakt heel snel fouten bij het omschrijven van machten en factoren 10 en milli-, kilo- en nano-dingen. De getalswaarde van de meting is door 1000 gedeeld en de onzekerheid door 100. (Fout 6)

$$g = 9,877 \pm 0,025 \text{ m/s}^2$$

Zucht. Het gaat de goede kant op, maar tussen een waarde en een eenheid hoort toch echt een *spatie*. (Fout 7) En je geeft onzekerheden op met één significant cijfer, tenzij dat cijfer een 1 is: dan geef je er twee. (Fout 8)

$$G = 9,87 \pm 0,02 \text{ m/s}^2$$

Soms heb je zo'n dag dat álles misgaat: beide afrondingen zijn verkeerd gedaan. (Fout 9 en 10) Bovendien staat er opeens een hoofdletter G, maar die betekent iets anders (Fout 11).

$$g = 9,88 \pm 0,03 \text{ m/s}^2$$

En nu zijn we er bijna! Deze notatie klopt wel. Er is alleen nog één probleem. Met een zo kleine onzekerheid is het verschil tussen de gevonden waarde en de literatuur meer dan twee keer de onzekerheid. Het “Zo dat is aardig gelukt” was een beetje te optimistisch. Er is gewoon een strijdigheid gevonden. En dus een verkeerde conclusie getrokken (Fout 12).

Het is wel extreem dat je de getalswaarde 9,87654321 zou vinden en het is ook extreem dat iemand *twaalf* fouten maakt bij het vinden van een onzekerheid. Maar misschien zijn er ook nog wel fouten die we niet gezien hebben. De slingerproef zelf kan verkeerd uitgevoerd zijn, waardoor de getalswaarde niet klopte. Er kan een rekenfout in hebben gezeten. De opgegeven waarde van de onzekerheid is wellicht te gunstig ingeschat.

Er moet wel nog een fout zijn, want de versnelling van de zwaartekracht ligt heus niet in het opgegeven bereik.

Twaalf fouten is wat overdreven, maar dit extreme voorbeeld laat wel zien hoe veel er mis kan gaan bij alleen al het noteren van de uitkomst en de onzekerheid. In dit boek noemen we de mogelijk foutoorzaken *De Grote Negen*.

1. eenheid

Vermeld de juiste SI-eenheden (en dus ook de correcte hoofdletters).

2. beste schatting

Als je geen waarde opgeeft, heb je geen meetresultaat.

3. onzekerheid

Laat de onzekerheid niet weg, anders kun je geen conclusies trekken.

4. **voorvoegsel**

Gebruik voorvoegsels als *milli* en *kilo* (met correcte hoofdletters).

5. **factor**

Ga correct om met machten van 10.

6. **significantie**

Geef de onzekerheid in het juiste aantal significante cijfers en de beste schatting met hetzelfde aantal decimalen.

7. **afronding**

Rond het cijfer 5 en hoger af naar boven en het cijfer 4 of lager naar beneden.

8. **haakjes**

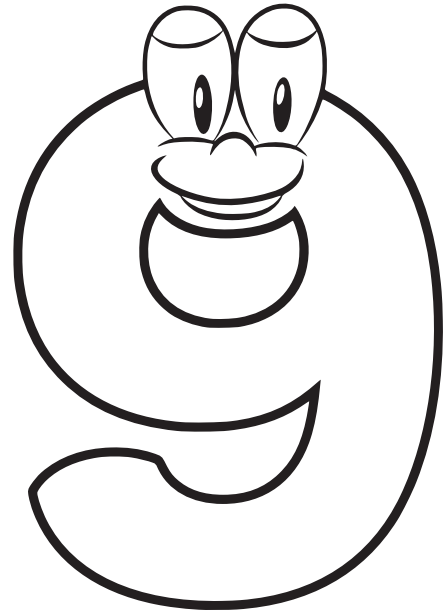
Zet haakjes om getallen als dat nodig is voor de rekenregels.

9. **komma**

Gebruik komma's in Nederlandse teksten (*and points in English*).

0. **spaties**

Zet een spatie tussen een getal en de eenheid. (Deze tiende van *De Grote Negen* heeft nummer 0, omdat nul niks is. Net als een spatie.)



Dit rijtje is zo belangrijk dat er nog wel een extra toelichting bij mag.

- De **eenheid** is het belangrijkste van alles. De lichtsterkte van een lamp druk je uit in lumen. Het vermogen ervan in watt. De massa in kilogram. Wie geen eenheid vermeldt, vertelt niet eens wat er gemeten wordt. “Een gloeilamp van 40 wát? Van 40 watt? Of van 40 lumen? Dat scheelt nogal wat...”
- De **beste schatting** en de **onzekerheid** erop zijn de getallen die aangeven wat je gemeten hebt. De kale cijfers dus, die allebei een keiharde onzekerheid hebben.
- Het **voorvoegsel** en de **factor** zijn twee verschillende dingen, maar ze hebben veel met elkaar te maken. Als je van 10^{-4} cm^3 naar 10^{-6} m^3 gaat, heb je met twee verschillende factoren (machten van 10) te maken, maar ook met een voorvoegsel ‘centi’ dat bij de ene term wel staat en bij de andere niet.
- De **significantie** in de beste schatting hangt af van de significantie in de onzekerheid. Niet *het aantal significante cijfers* is gelijk, maar *het aantal decimalen*.
- De **afronding** komt pas na het vaststellen van de significantie en kan zowel bij de beste schatting als bij de onzekerheid fout zijn: de beste

schatting kan verkeerd afgerond zijn en de onzekerheid ook.
(Eigenlijk twee mogelijke fouten... De Grote Negen is een tiental.)

- **Haakjes** hoeft je niet te zetten om het getallenpaar dat de beste schatting en de onzekerheid aanduidt, maar zijn wel belangrijk als je met factoren 10 werkt.
- Een **komma** in een Engelse tekst zetten of een punt in een Nederlandse: het is allebei onjuist als je het decimale scheidingsteken bedoelde. (Deze staat wat laag in het lijstje, omdat de fout bewust of onbewust door de lezer genegeerd wordt.)
- De **spatie** is heus wel verplicht tussen getal en eenheid, maar net niet afkeurenswaardig genoeg om hem tot *De Grote Negen* te rekenen. Als ik het heb over een afstand van 1609km, dan weet iedereen wat ik bedoel. Bij 1.609 km ligt dat anders.
What do you mean? One mile or a thousand miles?

Oefeningen

Een nieuwe meting komt uit op een onafgeronde waarde van 9,84452 m/s² en een onzekerheid van 0,17622 m/s². Noem de fouten in de notaties hieronder.

1. $9,8 \pm 0,2 \text{ m/s}^2$
2. $(9,84 \pm 0,17) \text{ m/s}^2$
3. $9,84 \pm 0,176 \text{ m/s}^2$
4. $9,85 \pm 0,18 \text{ m/s}^2$
5. $9.84 \pm 0.18 \text{ m/s}^2$
6. $(9,8 \pm 0,18) \text{ m/s}^2$
7. 9.844
8. Leid uit onderstaande formule af (zonder ergens op te zoeken dus!) wat de eenheid is voor de Gravitatieconstante G.

$$F = G \frac{m_1 \times m_2}{r^2}$$

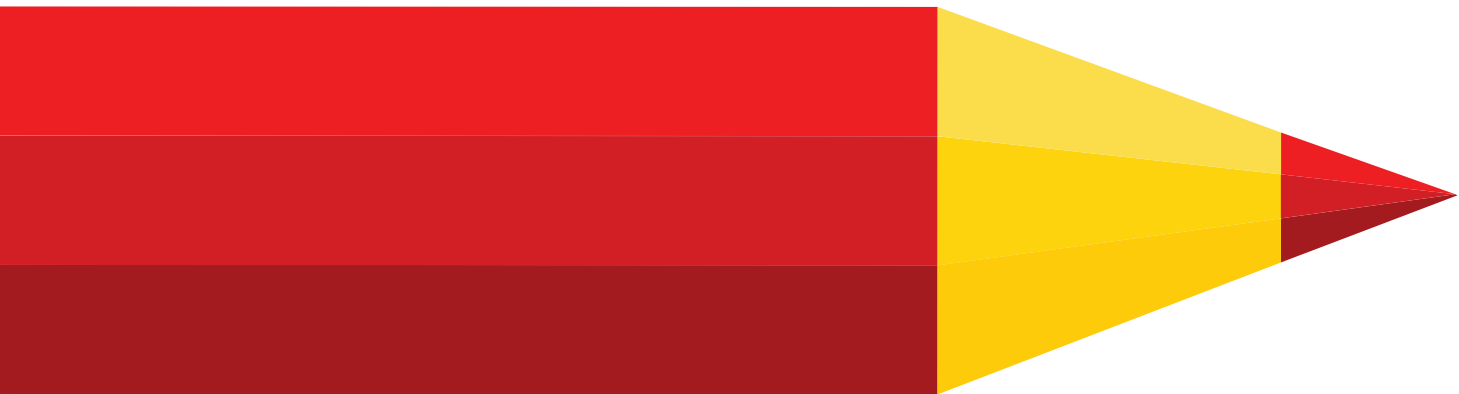
32. Verschillen meten

Als je de lengte van dit potlood zou willen meten, hoe doe je dat dan heel nauwkeurig?



Met een liniaal natuurlijk. Maar dan ben je er nog niet. Het potlood zal ook recht moeten liggen, en dit potlood ligt een graad of zes schuin. Je kunt zien dat het meer dan acht centimeter lang is... Maar hoe lang dan precies?

Laten we hem recht leggen en er dan eens flink op inzoomen.



Wow! Dat is nog eens inzoomen zeg! Het potlood blijkt toch een stukje langer te zijn dan ik dacht. Maakt het dan zoveel uit, dat hij recht ligt?

Maar wacht eens even...

Als we nauwkeurig willen meten, zullen we ook moeten controleren of de 0 exact bij het andere uiteinde van het potlood ligt. Dat kunnen we helemaal niet zien op deze twee afbeeldingen.

We kunnen niet alleen niet zien of het potlood exact bij de nul ligt, maar we kunnen ook niet zien of hij er een dikke centimeter naast ligt. Alles kan, zolang we het niet kunnen controleren.

Er kunnen twee dingen aan de hand zijn.

- Het potlood ligt zo goed mogelijk bij de nul, maar natuurlijk nooit helemaal exact. We kunnen daar een half streepje naast zitten. In dat geval hebben we te maken met een *onzekerheid* in de meting als gevolg van een onvermijdelijke fout.
- Het potlood ligt helemaal niet bij de nul, als gevolg van onoplettendheid of opzet. We kunnen er dan een flink stuk naast zitten. In dat geval hebben we te maken met een *verkeerde meting* in de meting als gevolg van al dan niet opzettelijke fout, die te vermijden was.

In dit vakgebied moeten we die twee dingen altijd goed uit elkaar houden.

- Je hebt fouten die een gevolg zijn van verkeerd meten.
- Je hebt fouten die zelfs optreden als je onberispelijk gemeten hebt.

Is die eerste categorie te vermijden? Nou ja, je ontkomt er als mens niet aan dat je fouten maakt. Dat is niet erg. Maar als je je best doet, maak je er minder. En als je het werk dat je zelf doet controleert én laat controleren, blijven er veel minder over.

Die tweede categorie, daar kom je niet onderuit. Wat je wél kunt doen is er rekening mee houden.

Als je de lengte van (bijvoorbeeld) een potlood opmeet, doe je eigenlijk *twee* metingen. Beter gezegd: je doet twee *aflezingen*. Aan beide uiteinden van het potlood lees je iets af. Aan de ene kant doe je je best om dat nul te laten zijn. Aan de andere kant schat je zo goed mogelijk tussen twee streepjes.

Een vuistregel is dat de onzekerheid gelijk is aan de helft van de afstand tussen twee streepjes. Omdat je dat twee keer hebt, stellen we de onzekerheid gelijk aan de hele afstand tussen twee streepjes. In dit geval: 1 mm.

33. Digitaal & Analooq



Figuur 33.1 Analoge klok

Hoe laat is het op bovenstaande klok? “Vijf voor twee.”

Bij een twaalfuursschaal zijn er dan nog twee mogelijkheden: het kan middag of nacht zijn. Op een digitale schaal: 13:55 of 01:55.

Deze klok heeft geen secondewijzer. Kunnen we het aantal seconden dan schatten? Dat ligt eraan. Bij sommige klokken verspringt de “minst significante” wijzer (minuten of seconden) met stappen en op andere klokken gebeurt dat in een vloeiende beweging. In het laatste geval kun je inderdaad schatten hoeveel seconden er voorbij zijn.

Deze klok heeft geen secondewijzer. Als een van beide wijzers afbreekt, dan mag je hopen dat het de minutenwijzer is. Die is het minst significant. Gek eigenlijk: de “grote wijzer” is minder belangrijk dan de “kleine wijzer”...

Hè, nou gebeurt het nog ook! Terwijl we nadenken over waarom ze de belangrijkste wijzer het kleinst gemaakt hebben, breekt die grote af.



Figuur 33.2 Klok met alleen een kleine wijzer

Kunnen we nu ook nog zien hoe laat het is?

“Tegen tweeën.”

De kleine wijzer heeft in ieder geval wel een vloeiend verloop. Hij staat duidelijk dichterbij de 2 dan bij de 1. Sterker nog: er staan vier streepjes tussen de 1 en de 2, dus dat gebied is in vijfen gedeeld. De afstand tussen twee strepen is $1/5$ uur oftewel 12 minuten.



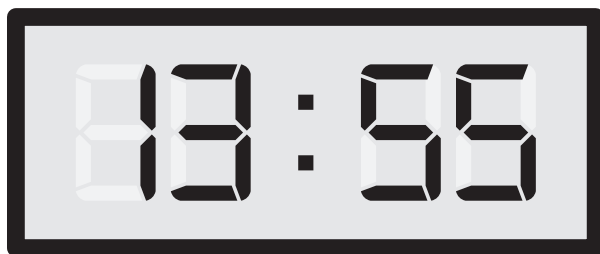
Figuur 33.3 Kleine wijzer, ingezoomd

Hè, was de kleine wijzer nou maar groter! Waren de klokkenbedenkers maar wijzer geweest. Dit is best nog lastig schatten. De wijzer staat dichterbij de dikke streep van de 2 dan bij de dunne streep van 12 voor 2...

Als ik zou moeten schatten, zeg ik: vijf voor twee.

En dat is nou juist zo lastig aan schatten. Het is bijna niet te doen om iets nog “eerlijk” te schatten als je al weet hoeveel het is. De uitslag van een korfbalwedstrijd kun je na afloop exact goed voorspellen. “Ik had wel verwacht dat PKC zou winnen. Achteraf.” Het psychologische aspect van *cognitieve bias* komt bij aflezen en meten ook voor.

Nu een andere klok.



Figuur 33.4 Digitale klok

Dat is andere koek: hier valt niks meer te schatten tussen streepjes.

Stel nu eens dat die analoge klok een minutenwijzer had die alleen maar precies op de streepjes kon staan. (Een wijzer die op z'n strepen staat dus.) Was het dan eigenlijk nog wel een analoge klok?

Kijk je in de *Van Dale Online*, dan vind je als betekenis van ‘analoog’:

niet werkend op basis van het binaire stelsel; gebruikmakend van wijzers bij het weergeven (tegenstelling: digitaal): een analoog horloge

Een klok met wijzers maakt gebruik van wijzers. *Wikipedia* zegt iets anders.

Analoog is de aanduiding voor een bepaald type technologie en de in sommige gevallen daarmee gepaard gaande signalen. Analoge meetinstrumenten geven een aanwijzing die in een continue relatie staat tot de te meten grootte.

Als de secondewijzer sprongetjes maakt, is dat dus niet het geval. De te meten grootte is de tijd die verstreken is sinds middernacht. Die kun je in seconden opgeven, maar ook in tienden ervan. Of in milliseconden. Of in nanoseconden of nog kleinere stapjes.

In dit hoofdstuk hebben we het over klokken, maar de grote lijn is van toepassing op analoge en digitale aflezingen in het algemeen. Daarbij moet gezegd worden dat je er bij een klok vanuit gaat dat de minutenwijzer niet “af-rondt”. Hij springt pas op vijf voor twee als het vijf voor twee is. Bij een digitale meter die iets anders meet dan de tijd, zou het kunnen zijn dat er een afronding plaatsvindt naar de dichtstbijzijnde waarde.

We hebben overigens één belangrijke stap overgeslagen. Wie zegt dat de klok gelijkloopt?

Oefeningen

1. Geef een schatting van de onzekerheid die optreedt bij het aflezen van de tijd op bovenstaande klok met één wijzer. Maak gebruik van zowel Figuur 33.1 als Figuur 33.2.
2. Hoe ziet de kansverdelingsfunctie eruit van tijd die je afleest op de digitale klok?
3. Welke eenheden staan er langs de assen van de grafiek van bovengenoemde kansdichtheidsfunctie?

34. Aflezen van een analoge schaal

Was het niet veel handiger geweest als die klok digitaal was?

15 : 02 : 48

Misschien. Als de seconden erbij gestaan hadden, zouden we in één keer hebben kunnen aflezen. Als de seconden er niet bij gestaan hadden trouwens ook, maar dan hadden we geen flauw idee gehad van dat gedeelte van de tijd. Als je 15:02 afleest, is het misschien wel één milliseconde voor 15:03.

Sommige analoge klokken hebben dat trouwens ook. De secondewijzer beweegt met sprongetjes van een minuut.

Dat laatste brengt ons op een interessante kwestie. Is een analoge klok in dat geval eigenlijk niet ook digitaal? Hij is analoog in de zin dat hij wijzers heeft. Hij is digitaal in de zin dat er alleen maar discrete waarden uit kunnen komen. Er valt niets af te lezen of te schatten tussen de minuten in.

In feite heeft ook de analoge klok met de continu roterende wijzer ook iets discreets in zich. Stel dat je een foto van zo'n klok maakt en je zou zo nauwkeurig mogelijk aflezen. Dan kom je misschien op 15:02:48,2. Dat betekent dat je twee tiende seconde afleest door heel goed te schatten. Misschien lukt het je zelfs met een extreem scherpe wijzer en een foto met hoge resolutie om te zeggen dat het 15:02:48,21 is of 15:02:48,214. Maar dan houdt het echt wel op. Ook een continue schaal wordt een keer discreet. Het aantal cijfers is nu eenmaal altijd eindig.

Hoeveel cijfers moet je zien af te lezen op een analoge of digitale schaal? Bij een digitale is het nogal makkelijk: alle cijfers die er maar staan. Er is geen goede reden om gegevens weg te laten die je hebt. Met de onzekerheid gaan we later pas rekening houden.

Bij een analoge meter zul je altijd moeten kiezen. Het aantal streepjes is eindig. Maar kun je ook nog iets zinnigs zeggen over wat er tussen de streepjes zit?



Figuur 34.1 Analoge meters voor spanning en stroom

Er is een vuistregel die zegt dat je tussen twee streepjes niet echt kunt aflezen, maar wel kunt zien welk streepje het dichtst in de buurt komt. Je kunt dus geen extra cijfer aflezen, maar je bent er wel zeker van dat je dichterbij het ene streepje zit dan bij het andere. De onzekerheid als gevolg van het aflezen zou dan een halve afstand tussen twee streepjes zijn.

Er valt wel wat af te dingen op die vuistregel. Wat als de wijzer van de schaal precies tussen twee streepjes lijkt te staan? Als je het dan niet helemaal goed ziet, zou je voor het streepje kunnen kiezen waar de wijzer juist iets verder vanaf staat. Zoals in het voorbeeld van de stroommeter hierboven, het rechter plaatje.

Aan de andere kant is het zo dat je bij een wat grove schaalindeling heus wel meer kunt zeggen dan bij welk streepje de wijzer het dichtst staat.



Figuur 34.2 Analoge spanningsmeters met verschillende schaalverdelingen

Neem de twee spanningsmeters hierboven.

De plaatjes zijn aan elkaar gelijk, het enige verschil is dat de kleine streepjes rechts weg zijn. Ja, je kunt zien dat de wijzer dichterbij de 400 V staat dan bij

de 300 V. Maar je kunt óók zien dat hij ergens halverwege ‘de tweede helft’ van dat gebied staat. Hij staat ergens tussen de 360 V en 390 V, laten we zeggen 375 ± 15 V. Een aanzienlijk kleinere onzekerheid dan 50 V dus.

Eigenlijk heb je gewoon niet zo veel aan die vuistregel. Kijk maar liever naar een onder- en bovengrens. Als je kunt schatten waar de werkelijke waarde tussen ligt, dan weet je het gebied en dus ook de onzekerheid.

Oefeningen

1. Leg uit waarom elke analoge meter in zekere zin ook digitaal is.
2. Waarvan hangt het af of je bij een analoge klok een schatting kunt maken van de tienden van de seconden?

35. Aflezen van een digitale schaal

Wat is de onzekerheid bij een digitale schaal?

Wat weet ik als het onderstaande getal op een digitale spanningsmeter staat, als die spanningen meet in volt?

483

Eigenlijk weet je niet wat je weet en niet weet, zolang dat niet in specificaties staat. Laat het apparaat de cijfers achter de komma gewoon maar weg? Of wordt er een afronding gedaan? Ligt de getalswaarde tussen 483,000... en 483,999...? Simpel gezegd: tussen 483 en 484? Of is er een 'nette' afronding gedaan en ligt de getalswaarde tussen 482,5 en 483,5?

In beide gevallen geldt: het komt sowieso wat vreemd uit met de significante cijfers. Want $483,5 \pm 0,5 \text{ V}$ en $483,0 \pm 0,5 \text{ V}$ suggereren beide een cijfer meer dan je eigenlijk hebt.

Om die reden geven we de onzekerheid bij digitale meters als gevolg van het aflezen meestal gewoon op als 1. De 483 op het display hierboven noteren we dus als $483 \pm 1 \text{ V}$. Het is een vuistregel; niet meer dan dat, maar ook niet minder.

Oefeningen

1. Hoe groot is de onzekerheid in de meting als je niet weet of de meter afrondt op gehele getallen of afkapt na het laatste cijfer?
2. Geef in de correcte notatie en met het juiste aantal significante cijfers de beste schatting en onzekerheid als het display 483 aangeeft en de eenheden *volt* zijn.

Foutdoorwerking

36. Foutdoorwerking bij optellen

Stel dat je de hoogte van een blikje wilt meten. Je doet dat met een duimstok en komt op 113,4 mm. Daar hoort een onzekerheid bij als gevolg van het aflezen. Laten we aannemen dat die onzekerheid een halve millimeter is. Wat is dan de hoogte van een stapel van zes blikjes?

Je zult eerst iets meer over de blikjes moeten weten. Bij frisdrank bijvoorbeeld zijn de blikjes zo gemaakt dat ze een beetje “in elkaar schuiven” als je ze opstapelt. Dat is handig, want dan vallen ze niet zo snel om. Maar voor het optellen van lengtes komt dat minder goed van pas.

Laten we ervan uitgaan dat we blikjes hebben waar dat niet het geval is. De hoogte van de stapel is dan $6 \times 113,4 = 680,4$ mm. Hoe groot is de onzekerheid daarin?

We hebben maar van één blikje de hoogte gemeten. De onzekerheid bedroeg 0,5 mm. Het zou dus kunnen dat de hoogte niet 113,4 mm was, maar 113,9 mm. Of 113,2 mm. Of 113,0 mm. Het zou ook 114,1 mm kunnen zijn, ook al is dat meer dan de beste schatting (113,4 mm) plus de onzekerheid (0,5 mm). Een onzekerheid geeft immers geen *harde grenzen* aan, maar een gebied waarin de werkelijke waarde *waarschijnlijk* ligt. Voor het rekenvoorbeeld in dit hoofdstuk doen we het even anders. We veronderstellen een uniforme verdeling tussen de genoemde grenzen en gaan uit van tienden van millimeters. Uitgedrukt in tienden van millimeters zijn er elf mogelijke ‘echte’ waarden: 112,9 mm, 113,0 mm, 113,1 mm, 113,2 mm ... 113,8 mm, 113,9 mm.

Mmmmm... Hierboven staat wel erg vaak een letter ‘m’. Hoe vaak eigenlijk? Dat ligt eraan wat je onder ‘hierboven’ verstaat.

Maar er is nog iets: we hebben ‘zomaar’ gezegd dat de onzekerheid 0,5 mm bedroeg. Hoe weten we dat? Wij als lezer weten het omdat het hierboven vermeld staat. Maar in praktijk zijn er metingen waar je over de onzekerheid zelf in het ongewisse verkeert. Verkeerde aannames liggen dan op de loer en dan klopt de rest van je berekening ook niet meer.

Stel dat we er 0,2 mm naast zaten. Dan zitten we er bij twee blikjes op elkaar $2 \times 0,2 \text{ mm} = 0,4 \text{ mm}$ naast en bij zes blikjes wordt dat $6 \times 0,2 \text{ mm} = 1,2 \text{ mm}$. Niet alleen de blikjes, maar ook de *fouten* stapelen zich op.

We hebben het nu over *fouten*, alsof we *weten* hoe groot het verschil is tussen de werkelijke waarde (*true value*) en onze beste schatting (*best estimate*),

wat in werkelijkheid natuurlijk niet zo is. Maar met een simpel rekensommetje willen we kijken of we inzien hoe zo'n fout doorwerkt bij het optellen.

De totale hoogte is volgens bovenstaande redenering maximaal gelijk aan $(680,4 \pm 3,0)$ mm. Dat is een significant cijfer te veel! Ook hier geldt: het gaat nu even over het snappen van hoe die fout doorwerkt in de optelling.

Wat nu als we twee blikjes met verschillende afmetingen hebben? Dan zijn er twee verschillende beste schattingen en twee verschillende onzekerheden. De redenering wordt daar niet anders door: de onzekerheden tel je bij elkaar op. Als het tweede blikje een hoogte heeft van $(162,7 \pm 0,5)$ mm, is de totale hoogte van beide blikjes minimaal $113,4 - 0,5 + 162,7 - 0,5 = 275,1$ mm en maximaal $113,4 + 0,5 + 162,7 + 0,5 = 277,1$ mm, wat je dus kunt schrijven als $(276,1 \pm 1,0)$ mm.

Zijn we nu niet wat te pessimistisch over de onzekerheid in die optelling? Voor het eerste blikje gold dat we een hoogte hadden geschat. Die hoogte kon ook iets groter of iets kleiner zijn, net als voor het tweede blikje. Het is dus ook mogelijk dat de werkelijke waarde (*true value*) van het eerste blikje 0,2 mm *groter* was dan wat we gemeten hebben en dat de waarde van het tweede blikje 0,2 mm *kleiner* was dan wat we gemeten hebben. Dan zouden de twee afwijkingen elkaar precies opheffen.

Voor dit moment volstaan we met het domweg optellen van de onzekerheden. In hoofdstuk 49 komen we terug op het feit dat fouten elkaar kunnen versterken, maar ook (gedeeltelijk) compenseren.

Voorlopige conclusie: als je de som van twee metingen berekent, moet je niet alleen de beste schattingen bij elkaar optellen, maar ook de onzekerheden.

Oefeningen

1. Twee weerstanden van 1,0 k Ω hebben een onzekerheid van 5%. Hoe groot (in procenten) is de relatieve onzekerheid van beide weerstanden samen als je ze in serie zet?
2. Aan 2,4 liter vloeistof (onzekerheid: 8%) wordt 600 cm³ vloeistof toegevoegd. Hoe groot mag de absolute onzekerheid in de toegevoegde vloeistof zijn om de relatieve onzekerheid in de totale vloeistof kleiner dan 10% te houden?
3. In dit hoofdstuk staat de vraag "Hoe vaak eigenlijk?" Bedenk een manier om die vraag te beantwoorden.
4. Is de vorige vraag niet raar voor een boek als *Metten & Weten*?

37. Doorwerking bij aftrekken

Bij het optellen van metingen met een onzekerheid moet je, uitgaande van het meest ongunstige geval, de onzekerheden ook optellen. Twee ijzeren staven hebben een lengte van een meter en een onzekerheid van 2 cm. Hoe lang zijn ze samen? Twee meter. Maar wat is de onzekerheid erop? In het meest ongunstige geval waren beide fouten ‘dezelfde kant op; en waren ze allebei 102 cm lang (of 98 cm) en zijn de dus samen 2,04 m of 1,96 m. Een verschil van maar liefst acht centimeter. Je zou de totale lengte moeten opgeven als $2,00 \pm 0,04$ m.

Hoe zit dat met het *verschil* in twee metingen? Stel, je gebruikt beide staven om een opstelling te maken en zet ze daarom naast elkaar. In het ideale geval is het verschil in hoogte exact gelijk aan nul. Ze zijn immers even lang. Hoe zit dat in het meest ongunstige geval? Die ene staaf kon weleens 102 cm zijn terwijl de andere misschien 98 cm is. Dan verschillen ze dus 4 cm in lengte. Als je het lengteverschil berekent en correct noteert, komt je op 0 ± 4 cm of $0,00 \pm 0,04$ m.

Je ziet dat beide antwoorden één significant cijfer hebben in de onzekerheid. Hoe groot is eigenlijk het aantal significant cijfers in $0,00$ m? Daar hebben we het in hoofdstuk 26 niet echt over gehad, maar het valt wel te beredeneren. $0,00$ is namelijk net zoiets als $0,01$ of $0,07$: het is een cijfer met twee nullen ervoor. Die laatste nul is dus een bijzonder geval: dat cijfer is wél significant.

Het zou ook een beetje vreemd zijn: een getal zonder significante cijfers heeft namelijk geen enkele betekenis.

Doordat we een verschil berekenden tussen waarden waarvan de beste schatting gelijk is aan nul, krijgen we de bijzondere situatie dat we een relatieve fout hebben die oneindig groot is. Beter gezegd: ongedefinieerd. De relatieve fout is namelijk $4 / 0$ en je kunt niet door 0 delen. Je zou ook kunnen zeggen dat je $0,00$ beschouwt als maximaal $0,005$. In dat geval kun je spreken van een relatieve fout van $4 / 0,005$. Dan heb je een factor 800, oftewel 80000 %. Niet oneindig, maar wel bizar groot.

Je ziet met dit eenvoudige voorbeeld dat wisselen tussen absoluut en relatief soms rare dingen oplevert. Het is voor een technicus wel belangrijk om telkens door beide glazen van dezelfde bril naar de dingen te kijken. Trouwens, niet alleen als technicus. Wie een winkeltje begint en op de eerste dag 14 artikelen verkoopt, moet niet zeggen dat hij of zij een oneindig grotere verkoop heeft dan de dag ervoor. Het is zelfs maar de vraag of het zo zinvol is om te

spreeken van ruim 50% stijging als je er de dag erna 22 verkoopt. Alleen grotere aantallen op langere termijnen zijn geschikt voor dergelijke analyses.

Terug naar de verschilmeting. Een verschil tussen twee positieve getallen is altijd kleiner dan de grootste van de twee. Maar de onzekerheid van het verschil is groter dan de grootste onzekerheid. Absolute fouten nemen toe bij het berekenen van een verschil. Relatieve fouten kunnen enorm groot worden, als de verschillen klein zijn.

Oefeningen

Twee weerstanden van $1,0\text{ k}\Omega$ hebben beide een onzekerheid van 5%. Het verschil tussen die weerstanden zou nul moeten zijn ($0\text{ }\Omega$), maar daar zit een onzekerheid op.

1. Hoe groot is die absolute onzekerheid?
2. Hoe komt het dat de relatieve onzekerheid in dit verschil veel groter is dan 10 %?

38. Een schuine tafel

Als we nu vier van zulke staven uit het vorige hoofdstuk gebruiken als poten van een tafel, kun je dan iets zeggen over hoe recht het tafelblad staat?

Laten we uitgaan van een onderlinge afstand tussen de poten van eveneens een meter. De grootst mogelijke helling krijg je dan als er aan de ene kant van de tafel poten zijn gebruikt van 102 cm en aan de andere kant poten van 98 cm. Het tafelblad heeft in dat geval een afloop van 4 cm per meter, oftewel 4%. De helling is dan te berekenen via de tangens... of toch de sinus? Is de overstaande zijde gedeeld door de aanliggende zijde gelijk aan 0,04? Of geldt dat voor de overstaande zijde gedeeld door de schuine zijde?

Als we rekenen met de tangens, komen we op een hoek van $2,29^\circ$ en als we kiezen voor met een sinus, blijkt het precies hetzelfde te zijn. Nou ja, precies: in twee decimalen. Het verschil is namelijk $0,00183^\circ$. Praktisch gezien niet relevant, maar welke van de twee neem je?

De vraag die hier achter zit, is of de tafel schuin staat of scheef is. Anders gezegd: staan de poten loodrecht onder het tafelblad en vormen ze een rechte hoek ermee, dan staan de poten zelf ook met een hoek van $2,29^\circ$ uit het lood. Maar als je de poten kaarsrecht overeind hebt staan, dan ligt alleen het tafelblad niet recht. Zoals gezegd: praktisch gezien geen verschil, maar je wilt voor je gaat rekenen wel een exacte definitie hebben van wat je gaat doen. Zodat iemand anders die het narekent met exact dezelfde gegevens exact dezelfde uitkomst krijgt.

Hebben we nu de helling van het tafelblad uitgerekend? Nee. Die $2,29^\circ$ is niet de helling, maar de *maximale* helling in het meest ongunstige geval. Eigenlijk hebben we de onzekerheid in de helling berekend. Omdat we vier even lange poten gebruikt hebben, zeggen we dat het blad onder een hoek van 0° staat. Maar omdat er een onzekerheid zit op de lengte van de poten, moeten we (in het juiste aantal significante cijfers) spreken van een hoek van $0^\circ \pm 2^\circ$.

We zien daarbij iets over het hoofd: de onderlinge afstand tussen de tafelpoten zal ook een onzekerheid hebben. Als die ook gelijk is aan twee centimeter, dan wordt de meest schuine stand bereikt als die vier centimeter afloop plaatsvindt over 98 cm. Ervan uitgaande dat de poten zelf ook scheef staan (en dus loodrecht onder het tafelblad), moeten we de inverse tangens uitrekenen van $0,04/0,98$. Dan kom je op een hoek van $2,34^\circ$.

Zijn we er nu? Nee. In de twijfel tussen sinus en tangens hebben we een keuze gemaakt voor poten die loodrecht onder het tafelblad staan. Maar ze

kunnen allebei een beetje scheef staan, of de ene kaarsrecht en de andere flink scheef.

Bij het nadenken over dit soort situaties kan het handig zijn om je af te vragen wat het meest extreme voorbeeld is. In dit geval: dat de ene poot kaarsrecht staat en de andere loodrecht erop. Dat zou wel erg raar zijn, maar in dat geval hebben we dus een tafelblad dat rechtop staat. Een hoek van 90° ... groter kan het niet worden, maar mogelijk is het wel.

Ondertussen hebben we het ook nog niet gehad over de dikte van de tafelpoot. Stel je een tafelpoot eens voor als een rechthoek die heel smal en hoog is. Wat gebeurt er als je een klein beetje naar rechts laat hellen? Dan kantelt hij wat, met als draaipunt de hoek rechts onderaan. De hoek links bovenaan komt juist wat omhoog. We zouden dus moeten weten hoe dik de poot is. En je voelt 'm al aankomen: ook in die dikte zit weer een onzekerheid.

Ook hier geldt: als je te veel in de details gaat en niet van ophouden weet, verlies je een gevoel voor de verhoudingen. We hebben tafelpoten van 1 m en daar zit een onzekerheid op van 2 cm. Als gevolg daarvan kent de hellingshoek van de tafel een onzekerheid van 2° .

Oefening

1. Een professionele waterpas heeft een tolerantie van 0,3 mm per meter. Bereken hoe groot het verschil tussen de lengtes van de poten van een tafel mag zijn om te zeggen dat hij 'waterpas' staat.

39. Doorwerken bij vermenigvuldigen

Zou je bij een vermenigvuldiging de fouten ook gewoon bij elkaar mogen optellen? Bij optellen was dat de oplossing, en bij aftrekken kwam dezelfde uitkomst naar voren... Het leven zou wel makkelijk zijn als je alleen maar hoefde op te tellen en dan de gevraagde uitkomst vond.

Het antwoord op de vraag is ‘nee’, en kan alleen al op basis van het nadenken over de dimensies gegeven worden. Een vermenigvuldiging zou namelijk uitgevoerd kunnen worden met twee grootheden met verschillende dimensies. Een moment bijvoorbeeld is het product van een kracht (in newtons) en een arm (in meters). En iets in newtons kun je niet optellen bij iets wat in meters wordt uitgedrukt.

Stel dat je een moment berekent. De kracht is $8,0 \pm 0,4$ N en de bijbehorende arm is 20 ± 2 cm. We hebben dan één relatieve onzekerheid van 5% en een van 10%. Wat is dan de onzekerheid in het product van die twee? Dat je de onzekerheden onderling moet vermenigvuldigen, zal ook wel niet... Dan kom je namelijk op 0,5% en dat is veel kleiner dan de kleinste van de twee.

Laten we eens uitrekenen wat er maximaal en minimaal uit de berekening kan komen. Of eigenlijk: de maximaal en minimaal *waarschijnlijke* waarde. De onzekerheid van 0,4 N garandeert niet dat je nooit boven de 8,4 N zult meten. Maar het klinkt wel erg redelijk om aan te nemen dat we de onzekerheden kunnen optellen bij en aftrekken van de beste schattingen om een marge te vinden waarbinnen het product ligt.

$$M_{min} = F_{min} \times r_{min} = 7,6 \times 0,18 = 1,368 \text{ Nm}$$

$$M_{max} = F_{max} \times r_{max} = 8,4 \times 0,22 = 1,848 \text{ Nm}$$

Als we het verschil tussen deze twee waarden door 2 delen, komen we uit op 0,240 Nm. Zonder rekening te houden met significante cijfers krijgen we dus $1,608 \pm 0,240$ Nm als uitkomst. De relatieve onzekerheid zou dan 14,93% zijn.

Is het toevallig dat we op de som van 5% en 10% uitkomen? Zou kunnen. Laten we het eens uitschrijven in symbolen. Dat is veel algemener en geeft vast een beter inzicht. En laten we ook met de groottes van kracht en arm rekenen en niet met de vectoren. (Dat deden we al niet.)

$$\begin{aligned} M_{max} = M + \Delta M &= F_{max} \cdot r_{max} = (F + \Delta F) \cdot (r + \Delta r) = \\ &= F \cdot r + F \cdot \Delta r + \Delta F \cdot r + \Delta F \cdot \Delta r \end{aligned}$$

Als we nu één slimme zet doen, zijn we bijna bij de oplossing. Deel links en rechts van het isteken beide door het moment. Rechts komt er dan het volgende:

$$\frac{M + \Delta M}{F \cdot r} = \frac{F \cdot r + F \cdot \Delta r + \Delta F \cdot r + \Delta F \cdot \Delta r}{F \cdot r}$$

Dat kun je ook schrijven als

$$1 + \frac{\Delta M}{M} = 1 + \frac{\Delta r}{r} + \frac{\Delta F}{F} + \frac{\Delta F \cdot \Delta r}{F \cdot r}$$

Aan beide zijden van het isteken kun je de 1 weglaten. Dat we na die ene slimme zet van de deling bijna bij de oplossing zijn, zit in die laatste term. Als we die zouden weglaten, dan stond hier dat de relatieve fout in M gelijk is aan de relatieve fout in r plus de relatieve fout in F .

Mogen we die laatste term dan zomaar weglaten? Nee, niet zomaar. Maar meestal wel. Onzekerheden zijn vaak (lang niet altijd!) veel kleiner dan de waarden waarop ze betrekking hebben. Daardoor is het product van twee relatieve fouten meestal heel klein. In ons voorbeeld: 5% van 10% is 0,5%. Door die term weg te laten, kom je op 15% in plaats van 15,5%. Daar ligt niemand wakker van.

Die 14,93% was overigens ook al niet gelijk aan 15,5% of 15%. Dat heeft te maken met het feit dat we niet met $F \cdot r$ zijn gaan rekenen als beste schatting: we gebruiken 1,608 N in plaats van 1,6 N.

We hadden bovenstaande formule trouwens ook kunnen uitschrijven voor de minimale uitkomst van het moment.

$$M_{min} = M - \Delta M = F_{min} \cdot r_{min} = (F - \Delta F) \cdot (r - \Delta r) = F \cdot r - F \cdot \Delta r - \Delta F \cdot r + \Delta F \cdot \Delta r$$

Ook dan kunnen we één slimme zet en links en rechts van het isteken door het moment delen.

$$\frac{M - \Delta M}{F \cdot r} = \frac{F \cdot r - F \cdot \Delta r - \Delta F \cdot r + \Delta F \cdot \Delta r}{F \cdot r}$$

Dat kun je ook schrijven als

$$1 - \frac{\Delta M}{M} = 1 - \frac{\Delta r}{r} - \frac{\Delta F}{F} + \frac{\Delta F \cdot \Delta r}{F \cdot r}$$

Dezelfde term kun je verwaarlozen en dezelfde enen kunnen schrappen. Je komt dan op dezelfde uitkomst, maar dan met overal minnen.

Wiskundig is het geen echt bewijs en het klopt ook niet exact. Maar de algemene formule is wel wat we gebruiken bij het doorrekenen van onzekerheden bij vermenigvuldigen.

Als $z = x \times y$

dan is

$$\frac{\delta z}{|z|} = \frac{\delta x}{|x|} + \frac{\delta y}{|y|}$$

We vinden dus óók bij vermenigvuldigen de onzekerheid in het product via optellen. We tellen nu alleen niet de *absolute* onzekerheden bij elkaar op, maar de *relatieve* onzekerheden.

Oefeningen

Een rechthoekig stalen vat heeft een hoogte van 1 m, een breedte van 4 m en een lengte van 8 m. Het is voor ongeveer de helft (je kunt daar nog 5% naast zitten) gevuld met water. De onzekerheden in de afmetingen van het vat bedragen 19 cm.

1. Welke onzekerheid levert de grootste bijdrage in de onzekerheid van het volume van het water?
2. Hoe groot zou het vat moeten zijn om ervoor te zorgen dat de onzekerheid in het hoogtepeil van het water evenveel invloed heeft op de onzekerheid in het volume als de afmetingen van het vat?
3. Bereken het volume van het water in het vat.
4. Het volume van een bol wordt gegeven door onderstaande formule.

$$V = \frac{4}{3}\pi r^3$$

Hoe groot mag de absolute onzekerheid in de straal van een bol met een diameter van 1,00 dm zijn om de onzekerheid in het volume minder dan 6% te laten zijn?

40. En als je het niet mag verwaarlozen?

Stel, we willen een vermogen berekenen door een spanning en een stroom te meten. Het vervelende is dat die spanning en die stroom beide nogal klein zijn en daardoor een grote relatieve fout hebben.

$$U = 0,2 \pm 0,1 \text{ V}$$

$$I = 5 \pm 3 \text{ A}$$

Uit je hoofd reken je uit dat het vermogen 1,0 W is. En dat de relatieve fouten 50% en 60% zijn. Bij vermenigvuldigen tel je de relatieve fouten bij elkaar op en kom je op 110%. Nog steeds uit het hoofd: de absolute fout is 1,1 W.

$$P = 1,0 \pm 1,1 \text{ W}$$

Huh? Mijn sommetje zegt dat ik misschien wel een negatief vermogen heb! En dat terwijl de waarden van de spanning en de stroom beide positief zijn.

Ze hadden trouwens ook beide negatief mogen zijn. Eén van beide negatief en de andere positief, kan natuurkundig gezien niet. Dan zou de stroom tegen stroom in zwemmen.

Een negatief vermogen, dat lijkt me sterk. En dat is het ook. Want die spanning is waarschijnlijk niet kleiner dan 0,1 V en de stroom is waarschijnlijk niet kleiner dan 2 A. Het vermogen zal dus niet onder 0,2 W liggen. De hoogste waarschijnlijke waarde vind je door 0,3 V te vermenigvuldigen met 8 A. Dat levert 2,4 W op. Een aardigere uitkomst zou dus zijn:

$$P = 1,3 \pm 1,1 \text{ W}$$

Het vervelende is alleen dat je beste schatting nu niet het product is van je beste schattingen. Dat is namelijk 1,0 W.

Waar komt nu toch dat negatieve vermogen vandaan? Hoe kun je ooit op -0,1 W uitkomen? Het antwoord is: die formules helpen je niet aan de exacte antwoorden, maar geven – zelfs als je corrigeert voor de term die we verwaarlozen – een grove schatting die alleen maar klopt voor kleine onzekerheden. (Dan klopt die schatting ook niet, maar maakt het niet zo veel uit.)

Als we in een tabel uitschrijven wat alle mogelijke uitkomsten zijn voor de combinaties van zeven verschillende stroomsterktes en vijf verschillende spanningen, dan vinden we deze waarden.

Tabel 40.1 Getalswaarden van het vermogen op basis van spanning en stroomsterkte

	2	3	4	5	6	7	8
0,10	0,20	0,30	0,40	0,50	0,60	0,70	0,80
0,15	0,30	0,45	0,60	0,75	0,90	1,05	1,20
0,20	0,40	0,60	0,80	1,00	1,20	1,40	1,60
0,25	0,50	0,75	1,00	1,25	1,50	1,75	2,00
0,30	0,60	0,90	1,20	1,50	1,80	2,10	2,40

In de linker kolom staan spanningen in volt (V). In de bovenste rij staan stroomsterktes in ampère (A). Alle andere velden zijn vermogens in watt (W). In het midden vind je keurig die 1,0 W die je zou verwachten. Maar die extreme waarden die uit de formule volgen, komen nergens voor.

Oefening

1. Bereken de gemiddelde waarde van de 35 mogelijke uitkomsten uit de tabel, en ga na of die gelijk is aan de 'middelste' waarde van 1,00. Geef een verklaring voor de overeenkomst of het verschil.
2. Lees de beide vragen hieronder door en beantwoord ze zonder eerst een berekening uit te voeren.
(Je mag niet rekenen, wel nadenken en schatten. Dat is iets anders.)
3. Bereken de gemiddelde waarde van het product van het aantal ogen dat je gooit met twee dobbelstenen.
4. Bereken de gemiddelde waarde van het kwadraat van het aantal ogen dat je gooit met één dobbelstenen.

41. Doorwerken in het algemeen

Niet alle grootheden zijn zo simpel dat ze te berekenen zijn via alleen optellen, aftrekken, vermenigvuldigen en delen. In dit hoofdstuk behandelen we een voorbeeld van een ‘lastigere’ functie. Het leuke daarvan is dat de aanpak zo algemeen is, dat hij op alle functies toepasbaar blijkt te zijn.

We nemen als voorbeeld de wet van Bragg, die opgesteld door de heer Bragg en zijn zoon. (Die zoon heette niet toevallig ook Bragg, dus eigenlijk is het de wet van Bragg & Bragg. We laten een van beiden weg, kijk zelf maar of je de vader of de zoon overhoudt.)

De wet van Bragg beschrijft hoe met diffractie van röntgenstralen de afstanden tussen de atomen in een kristalrooster kunnen worden berekend. Het gedeelte van de röntgenstralen dat wordt gereflecteerd, vertoont interferentiepatronen waaruit de afstanden kunnen worden afgeleid. De formule daarvoor luidt als volgt.

$$n\lambda = 2d \cdot \sin\theta$$

Hierin is

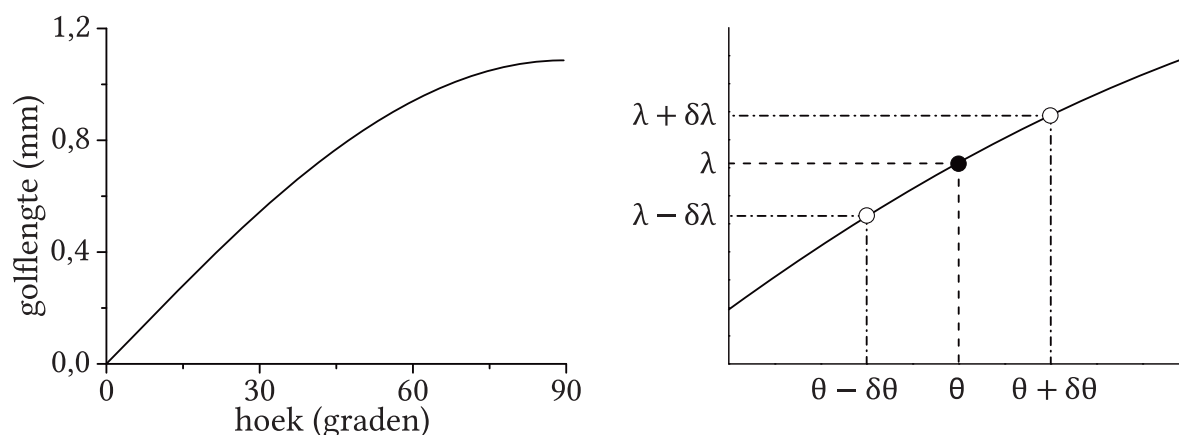
- n een geheel getal dat de diffractie-orde aangeeft;
- λ de golflengte van de röntgenstraling;
- d de afstand tussen de vlakken in het kristalrooster;
- θ de diffractiehoek tussen de stralen en de vlakken.

In dit voorbeeld gaan we uit van een eerste orde, dus $n=1$.

Als er in bovenstaande formule een grootheid s had gestaan waar nu $\sin\theta$ staat, dan hadden we te maken met een eenvoudige vermenigvuldiging van 2, d en s . Dan kon je gewoon de relatieve fouten optellen in die drie factoren. (Eigenlijk alleen in de laatste twee, want een factor 2 heeft geen onzekerheid.)

Als we aannemen dat de onzekerheid in d verwaarloosbaar klein is, dan houden we een functie over waaruit blijkt dat de onzekerheid in λ alleen maar afhangt van de onzekerheid in θ . Het lastige is: daar zit nog een sinus in. Hoe pakken we dat aan?

Kijk eens naar onderstaande twee figuren. Links staat een grafiek waarin de golflengte is uitgezet tegen de hoek. Rechts staat een klein stukje van diezelfde grafiek, sterk ingezoomd op het punt waarin we geïnteresseerd zijn.

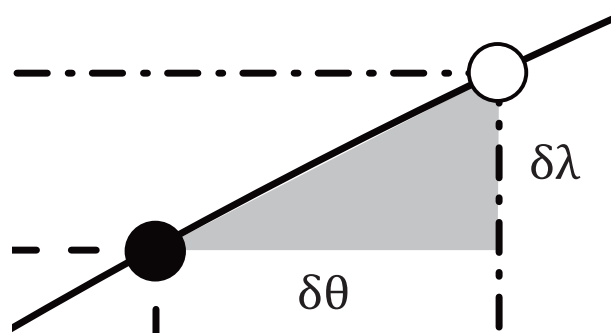


Figuur 41.1 Golflengte als functie van de hoek

Omdat we ver ingezoomd zijn, is het getoonde gedeelte van de grafiek bijna een rechte lijn geworden. Een rechte lijn, die dus te beschrijven is met een eenvoudige vergelijking:

$$\lambda = a\theta + b$$

Hoe groot is dan a en hoe groot is dan b ? Om met de laatste te beginnen: die interesseert ons hier niet. Laten we nog verder inzoomen op het plaatje.



Figuur 41.2 Golflengte als functie van de hoek (ingezoomd)

De lijn tussen de zwarte en de witte bol is net niet helemaal recht, maar het scheelt niet veel. Als we de lijn als recht beschouwen, dan geldt dat

$$\delta\lambda = |a|\delta\theta$$

De waarde van a hebben we tussen absoluutstrepen gezet, zodat het verhaal ook opgaat voor een dalende lijn.

We zouden hetzelfde kunnen doen voor de andere kant van de zwarte bol. Als de lijn echt recht is, maakt dat niet uit. Als er een kromming in zit die niet symmetrisch is, kom je op andere waarden uit.

Hoe groot is de waarde van a ? Die is gelijk aan de richtingscoëfficiënt van de raaklijn waarin we geïnteresseerd zijn. Die waarde kun je vinden door de afgeleide te bepalen in dat punt. We moeten dan de afgeleide bepalen van de functie, in ons geval van

$$f(\theta) = 2d \cdot \sin\theta$$

Die afgeleide is gelijk aan

$$\frac{df}{d\theta}(\theta) = 2d \cdot \cos\theta$$

In het punt $\theta = 60^\circ$ is de cosinus gelijk aan 0,5. Daar is de onzekerheid in de golflengte dus

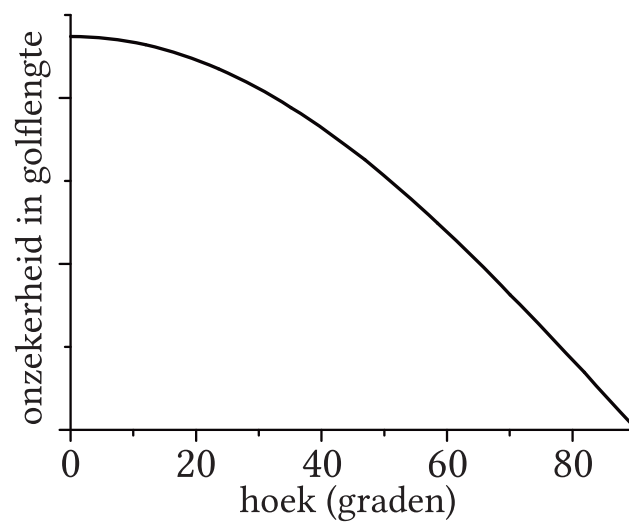
$$\delta\lambda = 2d \cdot \cos\theta \cdot \delta\theta = d \cdot \delta\theta$$

Zoals bij elke formule vragen we ons af hoe dat zit met de dimensies. De golflengte is in meters (of millimeters of nanometers, maar in ieder geval is het een lengte). De afstand d tussen de vlakken in het kristalrooster is ook een lengte. Maar $\delta\theta$ is toch in graden? Nee! Dit soort berekeningen kun je alleen maar uitvoeren als je in radialen rekent. Dan ben je namelijk af van de tamelijk willekeurige keuze dan een rechte hoek 90° is. Een rechte hoek is een kwart van een hele cirkel, die een omtrek heeft van 2π maal de straal. Rekenen in radialen geeft het enige juiste verband tussen de lengte langs de rand van de cirkel, de hoek en de straal.

Wat gebeurt er in de uiterste gevallen? Die sinus kan (in absolute waarde) variëren van 0 tot 1. Dat betekent dat de waarde van $\delta\lambda$, afhankelijk van de hoek maximaal gelijk is aan $2d \cdot \delta\theta$, en minimaal gelijk aan... nul! Dat kan toch niet? De onzekerheid in iets wat afhankelijk is van een grootheid waar onzekerheid in zit, kan toch niet nul zijn? Inderdaad. De afgeleide van de sinus (de cosinus dus) is *alleen maar* exact gelijk aan 0 bij een hoek van 90° (en 270°). De raaklijn loopt in dat punt precies horizontaal. Dat kun je ook zien in de grafiek die we eerder dit hoofdstuk bekeken. Vlak voor (en na) dat punt is de afgeleide ook erg klein, maar uiteraard niet nul.

In dit voorbeeld met de wet van Braggs geldt dus dat de onzekerheid voor kleine hoeken sterker afhangt van de onzekerheid in de die hoek dan voor hoeken van bijna 90° . De onzekerheid is geen constante waarde (absoluut) en ook niet een bepaald percentage (relatief), maar een functie van de hoek zelf.

Dat is weergegeven in de grafiek hieronder. Langs de verticale as zijn de waarden weggelaten, omdat die nog afhangen van $\delta\theta$. Het gaat om de vorm van de grafiek.



Figuur 41.3 Onzekerheid in de golflengte als functie van de hoek

42. Doorwerken bij delen

Voordat we kijken naar doorwerken bij delen, doen we het sommetje met vermenigvuldigen nogmaals, maar dan via een andere aanpak, namelijk de algemene benadering van het partiële differentiëren.

Als

$$z = x \cdot y$$

dan is

$$\delta z = \frac{\partial z}{\partial x} \cdot \delta x + \frac{\partial z}{\partial y} \cdot \delta y$$

In woorden: de onzekerheid in z is gelijk aan de som van de onzekerheden in x en y , vermenigvuldigd met hun partiële afgeleiden. Die zijn vrij eenvoudig bij een vermenigvuldiging, en dus kom je tot

$$\delta z = y \cdot \delta x + x \cdot \delta y$$

Delen we links en rechts door z , dan hebben we een uitdrukking voor de relatieve onzekerheid.

$$\frac{\delta z}{z} = \frac{y}{z} \cdot \delta x + \frac{x}{z} \cdot \delta y$$

Dan zijn we er al bijna, want

$$\frac{y}{z} = \frac{1}{x} \text{ en } \frac{x}{z} = \frac{1}{y}$$

Er staat dus precies wat we al wisten:

$$\frac{\delta z}{|z|} = \frac{\delta x}{|x|} + \frac{\delta y}{|y|}$$

De absolute-waardestrepen mag je toevoegen omdat onzekerheden altijd positief zijn.

Diezelfde benadering met partiële afgeleiden kun je toepassen op het quotiënt (de deling) van twee grootheden. Dan hebben we te maken met het volgende uitgangspunt.

$$z = \frac{x}{y}$$

De doorwerking van de onzekerheid berekenen we via

$$\delta z = \frac{\partial z}{\partial x} \cdot \delta x + \frac{\partial z}{\partial y} \cdot \delta y$$

wat resulteert in

$$\delta z = \frac{1}{y} \cdot \delta x + \frac{x}{y^2} \cdot \delta y$$

Omdat

$$y = \frac{x}{z}$$

geldt dus

$$\delta z = \frac{z}{x} \cdot \delta x + \frac{1}{y} z \cdot \delta y$$

Links en rechts delen door z leidt tot

$$\frac{\delta z}{|z|} = \frac{\delta x}{|x|} + \frac{\delta y}{|y|}$$

Zoals de formules voor optellen en aftrekken aan elkaar gelijk bleken te zijn, zijn de formules voor vermenigvuldigen en delen dat ook.

Te onthouden:

- Bij optellen (of aftrekken) van twee grootheden wordt de *absolute* onzekerheid in de som (of het verschil) gevonden door de *absolute* onzekerheden in beide grootheden bij elkaar op te tellen.
- Bij vermenigvuldigen (of delen) van twee grootheden wordt de *relatieve* onzekerheid in het product (of het quotiënt) door de *relatieve* onzekerheden bij elkaar op te tellen.

43. Doorwerken bij een exacte factor

Hoe zou de formule eruitzien voor de doorwerking van fouten als je een grootheid vermenigvuldigt met een exact getal?

Formule:

$$z = Bx$$

Eigenlijk is deze functie een bijzonder geval van een vermenigvuldiging.

Stel nu eens dat B niet een exact getal was. Dan had B een onzekerheid. Dan zou de relatieve onzekerheid in z gelijk zijn aan de som van de relatieve onzekerheden in B en x . Als B echter een exact getal is zonder onzekerheid, dan is die bijdrage nul. Conclusie: de relatieve onzekerheid in z is dan gelijk aan die in x . En aangezien z precies B keer zo groot is als x , zal δz precies B keer zo groot zijn als δx . Best wel logisch eigenlijk: als z én δz beide B keer zo groot zijn, is het quotiënt van die twee hetzelfde.

Eén ding moeten we nog in de gaten houden: er horen in de formule wel absolute strepen om B heen. Een relatieve en absolute fout zijn namelijk allebei altijd positief, ook als B negatief zou zijn.

Oefeningen

1. Toon voor het geval $B = 2$ aan dat de formule voor vermenigvuldigen met een exact getal consistent is met de formule voor doorwerken bij optellen.

44. Doorwerken bij machtsverheffen

Dit wordt weer een lekker kort hoofdstuk. Want wat is machtsverheffen méér dan gewoon vermenigvuldigen? Het kwadraat van een getal is dat getal keer zichzelf. Vermenigvuldig je de uitkomst nog een keer met dat getal, dan heb je de derde macht. Als je N dezelfde getallen met elkaar vermenigvuldigt, heb je de N^e macht.

$$z = x^2 = x \times x$$

$$\frac{\delta z}{|z|} = \frac{\delta x}{|x|} + \frac{\delta x}{|x|} = 2 \frac{\delta x}{|x|}$$

$$z = x^n = x \times x \times \dots \times x$$

$$\frac{\delta z}{|z|} = \frac{\delta x}{|x|} + \frac{\delta x}{|x|} + \dots + \frac{\delta x}{|x|} = |n| \frac{\delta x}{|x|}$$

Je ziet dat $|n|$ tussen absoluutstrepen staat. Een negatieve macht betekent een deling, en voor delingen geldt het dezelfde verhaal.

Oefening

1. Beredeneer of deze formule ook geldt voor $n = 0$.
2. Welke absolute onzekerheid is er groter: de absolute onzekerheid in het volume van een bol met een straal van x , of de absolute onzekerheid in het volume van een kubus met een zijde van x ?
3. Beantwoord de vorige vraag opnieuw, maar nu voor een bol met een *diameter* van x en een kubus met een zijde van x .

45. Doorwerken bij Louis Lyons

Lyons (1991) begint zijn eerste hoofdstuk met een eenvoudig voorbeeld van een slingerproef. Hij voert een proefpersoon op die op het volgende resultaat komt.

$$g = (9,81 \pm 0,15) \text{ m/s}^2$$

Je ziet dat Lyons beide getallen tussen haakjes zet. De meeste boeken doen dat niet, terwijl het eigenlijk correcter is. Zonder haakjes...

$$g = 9,81 \pm 0,15 \text{ m/s}^2$$

... staat er eigenlijk dat je een versnelling optelt bij een dimensieloos getal. Als een grootheid een getalswaarde is die je vermenigvuldigt met een eenheid, dan zouden er altijd haken moeten staan. Het punt is: dat is het niet. Je kunt een grootheid niet delen door een eenheid zoals m/s^2 . Wel door 1 m/s^2 .

Lyons analyseert vervolgens de uitkomst.

Possibility 1

If as suggested above the experimental error is ± 0.15 , then our determination looks satisfactorily in agreement with the expected value, i.e. 9.70 ± 0.15 is consistent with 9.81.

Wat je precies verstaat onder die onzekerheid, is altijd nog de vraag. Maar het moge duidelijk zijn dat 9,81 tussen de 9,65 en 9,85 ligt en dus niet strijdig is. Als de onzekerheid $0,10 \text{ m/s}^2$ geweest was, zou het volgens sommige definities nog steeds niet strijdig zijn geweest. Daarna zegt Lyons iets vreemds.

Possibility 2

If we had performed a much more accurate experiment and had succeeded in reducing the experimental error to ± 0.01 , then our measurement is inconsistent with the previous value. Hence, we should worry whether our experimental result and/or the error estimate are wrong. Alternatively, we may have made a world shattering discovery, i.e. 9.70 ± 0.01 is inconsistent with 9.81.

Dat laatste is wetenschapsfilosofisch interessant. Wat doe je als je vaststelt dat een algemeen aanvaarde waarde niet klopt? De essentie van wetenschap is dat het controleerbaar is. Veel mensen zullen de literatuur waarde voor waar aannemen en de fout bij zichzelf zoeken. Op zich een prijzenswaardig voorbeeld van bescheidenheid én realiteitszin. Maar als je altijd aan jezelf gaat twijfelen als je iets anders vindt, hoe kun je de dingen dan nog toetsen?

Waar ik over val, is de volgende passage uit de tekst van Lyons hierboven.

If we had performed a much more accurate experiment and had succeeded in reducing the experimental error to ± 0.01 , then our measurement is inconsistent with the previous value.

Dit is moeilijk vol te houden. Als we een nauwkeuriger experiment hadden uitgevoerd, zouden er namelijk – als we het goed deden – twee dingen veranderen. De onzekerheid zou inderdaad kleiner geworden zijn. Stel nu eens dat die van $0,15 \text{ m/s}^2$ zou zijn afgenomen naar $0,02 \text{ m/s}^2$, zouden we dan wéér $9,70 \text{ m/s}^2$ hebben gevonden? Als dat waar is, vraag ik me af wat er zo nauwkeurig aan dat nieuwe experiment was. Wie de valversnelling met een onzekerheid van $0,02 \text{ m/s}^2$ bepaalt, weet één ding zeker. Die meet iets van tussen pakweg $9,79$ en $9,82 \text{ m/s}^2$. Het kan pakweg twee keer de onzekerheid afwijken van de werkelijke waarde, maar onder de $9,75 \text{ m/s}^2$ kom je domweg nooit, omdat je nu eenmaal nauwkeurig gemeten hebt.

Je zou dus een kleinere onzekerheid gevonden hebben én een ‘betere’ beste schatting.

Possibility 3

If we had been stupid enough to time only one swing of the pendulum, then the error on g could have been as large as ± 5 . Our result is now consistent with expectation, hut the accuracy is so low that it would be incapable of detecting even quite significant differences. i.e. 9.70 ± 5 is consistent with 9.81 , and with many other values too.

Hier moet ook een kanttekening bij worden geplaatst. Iemand die een fout maakt, is niet *stupid*. Je bent niet *stom* als je iets verkeerd doet. Je hebt gewoon iets verkeerd gedaan en als je daarachter komt, leer je en heb je de kans om het beter te doen. Wie iets deed wat hij of zij anders had moeten doen, had waarschijnlijk een onjuiste gedachte. En als je gedachte niet klopt, ben je zelf nog niet fout. Je bent immers niet je eigen gedachte. (En ook niet die van een ander.)

Mag je zeggen dat het *een domme fout* was? Dat niet degene die het deed stom of dom is, maar *de fout*? Tsja... Als er domme fouten bestaan, dan zijn er waarschijnlijk ook slimme fouten... Die kom je toch maar zelden tegen.

Nu inhoudelijk. Dat de onzekerheid maar liefst 5 m/s^2 zou kunnen worden door slechts één keer te slingeren, leek mij sterk toen ik het voor het eerst las. Dat was met een goed gekoeld glas met iets lekkers erbij, buiten in de zon, na werktijd en aan het einde van een schooljaar. Toen ik mijn ogen van het boek ophief, dacht ik: ‘Wat een onzin!’

Laten we er eens aan rekenen, want dat kunnen we inmiddels.

De formule voor de slingertijd (voor kleine hoeken) luidt als volgt.

$$T_0 = 2\pi \sqrt{\frac{l}{g}}$$

Uiteraard schrijven we die eerst zo dat g wordt uitgedrukt in de andere grootheden.

$$g = \frac{4\pi^2 l}{T_0^2}$$

Het getal π kennen we in zoveel decimalen dat die term weggelaten mag worden. De tijd zit kwadratisch in de formule. We hebben dus dit over.

$$\frac{\delta g}{|g|} = \frac{\delta l}{|l|} + 2 \frac{\delta T_0}{|T_0|}$$

Als nu de onzekerheid door een andere manier van tijdmeten toeneemt van 15 % naar iets meer dan 50 %, dan zou dat betekenen dat de onzekerheid in de tijd de helft van het verschil tussen die twee percentages bedraagt. Dan gaat het dus om 18%. Het is niet heel vreemd om 0,3 s onzekerheid in je tijdsmeting te hebben, dus bij een slingertijd van 1,7 s gaat het inderdaad zo hard mis.

Een formule kan helpen om de lengte van de slinger te vinden.

$$l = \frac{g T_0^2}{4\pi^2}$$

Altijd als ik zo'n formule zie, kan ik het niet laten om de eenheden te controleren: g is in m/s^2 en T_0 is in s. Aangezien T_0 in het kwadraat gaat en de noemer dimensieloos is, komen er inderdaad meters uit. Gelukkig maar.

De slinger zou 72 cm zijn. Aan de korte kant, maar het zou goed kunnen. Die uitspraak 'Wat een onzin!' neem ik terug.

Eén ding dat Louis Lyons daarna nog zegt, staat buiten kijf.

The moral is clear. Whenever you determine a parameter, estimate the error or your experiment is useless.

Afhankelijk & Onafhankelijk

46. Voorbeelden van (on)afhankelijkheid

Voordat het we het over ‘gewoon’ en kwadratisch optellen gaan hebben in deze sectie, bekijken we eerst twee voorbeelden die duidelijk maken wat we eigenlijk met *afhankelijk* bedoelen.

We willen de oppervlakte van een tafelblad berekenen en gebruiken een rolmaat om de lengte en de breedte te bepalen. De oppervlakte volgt uit de formule

$$A = l \times b$$

Bij een vermenigvuldiging tel je de relatieve fouten bij elkaar op. Dus:

$$\frac{\delta A}{A} = \frac{\delta l}{l} + \frac{\delta b}{b}$$

De absoluutstrepen zijn hier weggelaten, omdat we met alleen maar positieve waarden te maken hebben.

Hoe zit het nu met de *afhankelijkheid* tussen l en b ? Wat als onze rolmaat systematisch 1% te veel aangeeft? Dat een tafelblad met een lengte van 120,0 cm door ons opgemeten wordt als 121,2 cm? Dat zal dan doorwerken in de oppervlakte. Volgens de formule hierboven is de bijdrage 1%.

Diezelfde rolmaat zal bij het opmeten van de breedte hetzelfde euvel vertonen. Ook daar is het 1% te veel. Het tafelblad van 80,0 cm wordt door ons opgemeten als 80,8 cm.

Voor de zekerheid meten we de lengte en de breedte nog twee keer. Geloof het of niet: we vinden nóg twee keer 121,2 cm en 80,8 cm. Is dat even toevallig? Nee. De rolmaat heeft een *systematische* fout. Die vind je niet door te herhalen. Hij is niet zichtbaar, en zou gevonden moeten worden door een kalibratie.

De werkelijke oppervlakte van de tafel zou je kunnen uitrekenen als je de werkelijke waarden van de lengte en de breedte wist.

$$A = l \times b = 1,200 \times 0,800 = 0,960 \text{ m}^2$$

Reken je met de door ons gemeten waarden, dat vind je net iets anders.

$$A = l \times b = 1,212 \times 0,808 = 0,979296 \text{ m}^2$$

Dat is natuurlijk een te groot aantal significante cijfers, maar we moeten de onzekerheid nog uitrekenen. Die weten we toch? Die was toch 1%? Nee! De

systematische fout is in ons zelfverzonnen voorbeeld gelijk aan 1%. In werkelijkheid weet je niet dat een rolmaat exact 1% te veel aangeeft. (Dan zou je er voor corrigeren. Bovendien: geen enkele fabrikant maakt een rolmaat die exact 1% te veel aangeeft. Als hij dat wist, zou hij de streepjes beter zetten.)

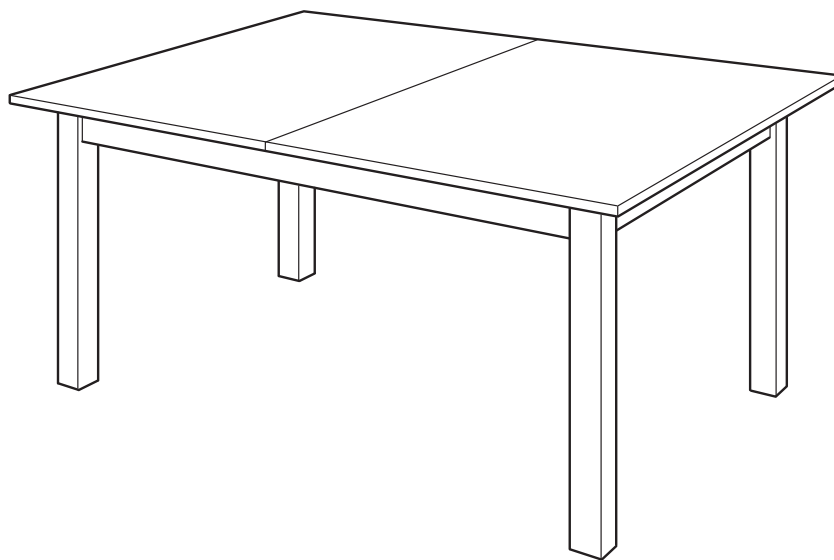
Laten we aannemen dat de fabrikant opgegeven heeft dat de rolmaat een onzekerheid heeft van 2%. Dan is de onzekerheid in de oppervlakte het dubbele daarvan. En 4% van 0,979296 is 0,03917184. De onzekerheid wordt opgegeven in één significant cijfer. Dat resulteert in de volgende uitkomst.

$$A = 0,98 \pm 0,04 \text{ m}^2$$

De uitkomst is niet strijdig met de werkelijke waarde.

Hoe zit het met de afhankelijkheid? Wat nu als we een rolmaat hadden gehad die precies 1% te weinig had aangegeven? Dan waren zowel de lengte als de breedte te klein geweest.

Het moge duidelijk zijn dat hier sprake is van een zeer sterke afhankelijkheid. Als de gemeten lengte te groot is, is de gemeten breedte ook te groot. En als de lengte te klein is, is de gemeten breedte ook te klein. Het meten met hetzelfde instrument geeft een sterke afhankelijkheid.



Figuur 46.1 Zomaar een tafel, om de lege ruimte op de pagina te vullen (afbeelding van IKEA)

Zomaar een tafel? Welnee. Dit is een BJÖRNA. Kenners zien dat zo.

Nu een ander voorbeeld. Stel dat je het rendement van een motor wilt berekenen. De motor tilt een blok met een gegeven massa op tot een bepaalde hoogte, in een bepaalde tijd. Op de motor staat een spanning en erdoorheen

loopt een stroom. Het rendement is de benodigde energie om het hoogteverschil te overbruggen gedeeld door het vermogen van de motor maal de tijd.

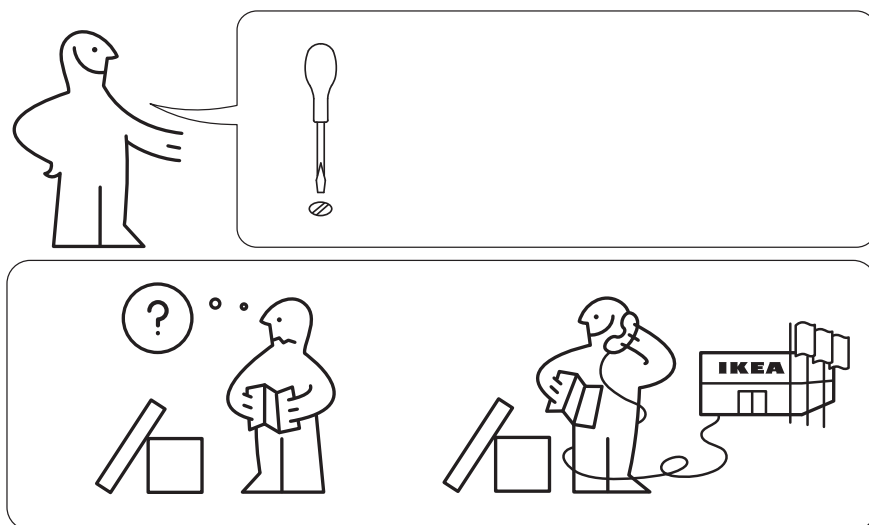
$$e = \frac{\text{geleverde arbeid}}{\text{verbruikte energie}} = \frac{mgh}{VIt}$$

Rechts van het isteken staan maar liefst zes grootheden waarvan er eentje wordt opgezocht in de literatuur (de gravitatieversnelling ter plaatse) en de andere vijf door vijf verschillende instrumenten wordt opgemeten. De relatieve fout in de uitkomst zou je kunnen berekenen door

$$\frac{\delta e}{e} = \frac{\delta m}{m} + \frac{\delta g}{g} + \frac{\delta h}{h} + \frac{\delta V}{V} + \frac{\delta I}{I} + \frac{\delta t}{t}.$$

Wat denk je, zal hier een te grote m ook zorgen voor een te grote g of h ? Of zal de fout in je stopwatch doorwerken op de multimeter? Nee. En dus is het domweg optellen van de relatieve onzekerheden veel te pessimistisch.

Van belang is dat je inziets dat die metingen bij dat tafelblad *afhankelijk* waren. De grootheden in het voorbeeld van de motor zijn *onafhankelijk*.



Dat mannetje van de IKEA-handleidingen maakt altijd een wat sukkelige indruk op me. Kijk eens naar die domme grijs als hij uitroept dat hij alleen maar een schroevendraaier nodig heeft. En dan komt hij er nóg niet uit. Nou ja, dan maar bellen naar de zaak waar het vandaan komt. Met een telefoon die kennelijk *in de winkel* staat. Volg het snoer maar.

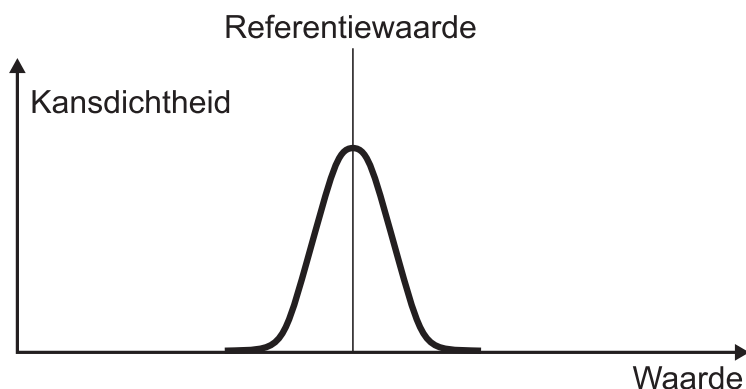
47. Extreme voorbeelden

Stel, wel hebben twee (fictieve) multimeters, die op het eerste gezicht best op elkaar lijken. Vooral aan de buitenkant. Zoek de verschillen!

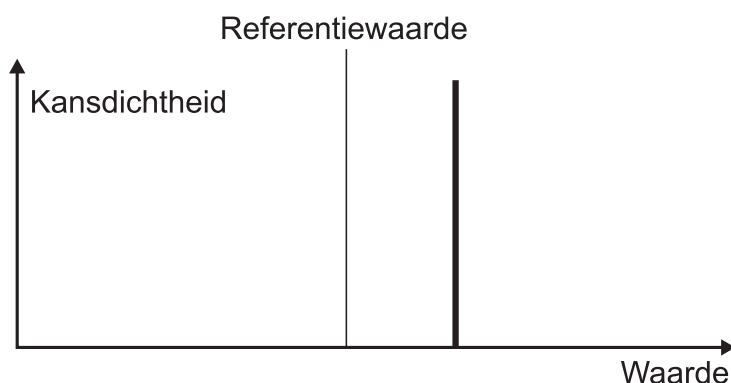
Er is maar één verschil:

- meter A (de bovenste) geeft meetwaarden die extreem nauwkeurig zijn (maar niet erg precies)
- meter P (de onderste) geeft meetwaarden die extreem precies zijn (maar niet erg nauwkeurig)

In het *Wikipedia*-plaatje waarin *accuracy* and *precision* worden geschetst, zien hun kansverdelingen er als volgt uit.



Figuur 47.1 Metingen van Accurate Digital Multimeter A



Figuur 47.2 Metingen van Precise Digital Multimeter P

Gemiddeld genomen geeft meter A (de nauwkeurigere, *accurate*) de juiste waarde, maar er zit spreiding in de uitkomsten. Meter



P, de *precieze* variant geeft bij elke herhaalde meting hetzelfde resultaat. Maar die meting wijkt wel een (onbekend) stuk af van de werkelijke waarde.

Beide apparaten hebben bij het meten van een weerstand een onzekerheid van $10\ \Omega$. Die komt op een verschillende manier tot uiting.

- meter A geeft normaal verdeelde metingen met een standaardafwijking van $10\ \Omega$ en een gemiddelde dat gelijk is aan de werkelijke waarde
- meter P geeft telkens dezelfde afwijking, die met een kans van ongeveer twee derde binnen $10\ \Omega$ van de werkelijke waarde ligt.

Voor meter P geldt dat de afwijking van de werkelijke waarde $10\ \Omega$ zou kunnen zijn, maar het kan ook best $8\ \Omega$ of $11\ \Omega$ zijn. Of $7\ \Omega$, of $9\ \Omega$, of $12\ \Omega$... Het moet in die orde van grootte liggen, maar je weet het niet exact. Anders zou je ervoor kunnen compenseren. (Dan is het geen *onzekerheid* meer.)

Voor meter A geldt dat je de onzekerheid net zo klein kun krijgen als je maar wilt, door veel metingen uit te voeren. Je bepaalt dan namelijk gewoon het gemiddelde van een groot aantal metingen. We weten dat de standaardafwijking van de weerstanden $\sigma_R = 10\ \Omega$. Het gemiddelde van een (groot) aantal metingen is op zichzelf ook normaal verdeeld, maar wel met een kleinere standaardafwijking, namelijk

$$\sigma_{\bar{R}} = \frac{\sigma_R}{\sqrt{N}}$$

Bij honderd metingen heb je dus in het gemiddelde maar een onzekerheid van $1,0\ \Omega$. Als de onzekerheid $0,10\ \Omega$ zou moeten zijn, moet je nóg honderd keer zo vaak meten: tienduizend keer.

Het is niet erg realistisch dat je twee multimeters zou hebben waarvan de ene geen systematische fout heeft en wel een toevallige en de andere juist wel een systematische en geen toevallige. Maar het is wel nuttig om er dit gedachtenexperiment mee te doen.

Stel dat je een (ook weer fictieve) weerstand zou hebben van exact $200\ \Omega$. Je meet die tien keer met Meter P (die een volmaakte precisie heeft) en vindt dan: $207\ \Omega$, $207\ \Omega$, $207\ \Omega$, $207\ \Omega$, $207\ \Omega$, $207\ \Omega$, $207\ \Omega$, $207\ \Omega$, $207\ \Omega$ en $207\ \Omega$. Ja, het is best een saaie Piet, die Meter P. Tien keer hetzelfde: je zou bijna gaan denken dat er geen onzekerheid op de meting zit. Maar die is er wel! De onzekerheid van $10\ \Omega$ betekent dat de werkelijke waarde waarschijnlijk (met een kans van ongeveer 2 op 3) ligt tussen $197\ \Omega$ en $217\ \Omega$.

Nu de andere multimeter, Meter A. Die vindt wél tien verschillende waarden. $188\ \Omega$, $195\ \Omega$, $202\ \Omega$, $226\ \Omega$, $202\ \Omega$, $200\ \Omega$, $186\ \Omega$, $196\ \Omega$, $185\ \Omega$ en $195\ \Omega$.

De Meter A is fictief en we willen wel “echte” normaal verdeelde metingen vinden. Het is lastig om die te bedenken. ‘Zomaar een getal onder de 10... Ergens op internet stond dat mensen meestal ‘zeven’ zeggen.

Hoe komen we nu aan tien willekeurige getallen die normaal verdeeld zijn met een gemiddelde van 200 en een standaardafwijking van 10? Daar hebben we *Excel* voor! Die heeft het bovenstaande rijtje verzonnen.

Nu we de beide reeksen hebben, komen we tot de kern van dit hoofdstuk.

- De metingen van Meter A zijn onderling *volstrekt onafhankelijk*.
- De metingen van Meter P zijn onderling *volstrekt afhankelijk*.

Dat heeft gevolgen voor het doorwerken van onzekerheden bij optellen.

Als we twee weerstanden van exact $200\ \Omega$ in serie zetten, dan is de vervangingsweerstand gelijk aan $400\ \Omega$.

Meter P zal voor de eerste weerstand zeggen dat die $207\ \Omega$ is... *en voor de tweede weerstand ook!* De meter heeft een afwijking die de gebruiker niet exact kent, maar we weten wel dat die constant is en (absoluut) in de orde van grootte van $10\ \Omega$ ligt. Maar de afwijking is altijd hetzelfde en versterkt dus zichzelf bij optellen.

Meter A zal voor de eerste weerstand zeggen dat die $188\ \Omega$ is en voor de tweede weerstand dat die $195\ \Omega$ is. Een eventuele derde is $202\ \Omega$ en ga zo het rijtje van willekeurige afwijkingen maar af. Telkens is de volgende weerstand gebaseerd op *toeval*. De waarde is gemiddeld gelijk aan $200\ \Omega$ en is de ene keer lager en de andere keer hoger. Het is dus wat pessimistisch om te zeggen dat die spreiding bij twee weerstanden twee keer zo groot is.

Wiskundig kun je bewijzen dat de som van twee normale verdelingen ook normaal verdeeld is. Het gemiddelde van die ‘opgetelde’ verdeling is gelijk aan het gemiddelde van beide gemiddelden. De variantie van de ‘opgetelde’ verdeling is gelijk aan de som van beide varianties.

In ons voorbeeld is de verwachtingswaarde van de meting voor de eerste en de tweede weerstand gelijk aan $200\ \Omega$. De variantie is voor beide gelijk aan $(10\ \Omega)^2 = 100\ \Omega^2$. De som is dus $200\ \Omega^2$ en de standaardafwijking is daar de wortel van: $10\sqrt{2}\ \Omega \approx 14,142135623731\ \Omega$. In twee significante cijfers: $14\ \Omega$.

48. “Een getal onder de tien?!”

Nog even over die normaal verdeelde toevallige fouten. Die zijn in *Excel* gegenereerd met deze formule.

```
=AFRONDEN(NORM.INV.N(ASELECT();200;10);0)
```

Een Engelstalige *Excel* had de volgende formule nodig gehad:

```
=ROUND(NORM.INV(RAND();200;10);0)
```

Het eerste argument moet een willekeurige kans zijn, een getal tussen 0 en 1. Dat is precies wat je met `ASELECT()` of `RAND()` doet. Het tweede argument is het gemiddelde en het derde argument is de standaardafwijking. De functies `AFRONDEN()` en `ROUND()` hebben de uitkomst van de functie als argument én het getal 0, omdat je geen decimalen wilt hebben.

Dit is een “dubbel geneste” functie. De uitkomst van `ASELECT()` is een argument van `NORM.INV.N()` en de uitkomst daarvan is een argument voor `AFRONDEN()`.

Nu de psychologie. Net zoals het lastig is om een écht willekeurig getal onder de tien te roepen, valt het ook niet mee om een willekeurig gegenereerd rijtje normaal verdeelde getallen als “willekeurig” te beschouwen. De neiging bestaat om *Excel* net zo lang een nieuw reeks te laten genereren tot ze er normaal genoeg en toevallig genoeg uitzien. Een te grote uitschieter bijvoorbeeld, dat voelt als abnormaal. Terwijl uitschieters de normaalste zaak van de wereld zijn. Ze komen minder vaak voor dan andere waarden, maar als ze er niet zijn, is er iets geks aan de hand.

Ik heb *Excel* een getal met 68 cijfers laten genereren. Dat kan niet, want met zoveel cijfers rekt hij niet. Dus nam ik er zeventien keer vier en plakte die achter elkaar. Dit is wat eruit kwam.

73527662437853989549300078032157063837626886550209286606194650233934

Je ziet dat er een setje van drie nullen in zit. En ook twee keer twee zessen. Komen die soms vaker voor dan de andere cijfers? Laten we het rijtje eens sorteren van klein naar groot.

00000000011222222233333333334444555555566666666667777778888888999999

Het lijkt toch mee te vallen, al zijn er wat weinig enen. Klopt dat wel?

Ik genereerde opnieuw (in twee stappen) een rijtje van 68 cijfers.

21332687989760391670265674957518614205496573644541028952078955703452

Geen setjes van drie keer dezelfde achter elkaar. Wel een keer twee drieën, en twee vieren, en twee vijven... Het vreemde is: juist doordat het toeval is, krijg je soms patronen die je zelf niet zomaar zou bedenken. De meeste mensen zullen als ze een rij willekeurige getallen onder de tien moeten opnoemen zelden twee keer dezelfde achter elkaar kiezen. En al helemaal niet drie keer dezelfde!

De kans dat een cijfer gevolgd wordt door hetzelfde cijfer, is 1 op 10. De kans dat dat *twee keer gebeurt*, is 1 op 100. Hoe groot is nou de kans dat iets wat in 99% van de gevallen niet gebeurt *68 keer achter elkaar* niet gebeurt?

$$P = (0,99)^{68} \approx 50\%$$

Deze kans is dus een half. Een setje van drie is net zo waarschijnlijk als geen setje van drie.

Eigenlijk moet je natuurlijk niet tot de macht 68 rekenen, maar tot de macht 66, want de laatste twee cijfers hebben geen twee cijfers meer achter zich. Dan komt je afgerond op 51,5%.

Waarom nam ik eigenlijk 68 cijfers? Zomaar? Nee. Het is het aantal cijfers dat op één regel past. Eerst had ik er 64 genomen (een macht van 2), maar toen kwam de kans die ik uitrekende niet op een half uit.

De gewetensvraag is: wat zou ik gedaan hebben als er bij het genereren van een setje cijfers géén drietal in zat? Dat had net zo goed gekund... En dan had ik het heus wel overgedaan. Zó moeilijk is het om iets toevalligs te doen.

Iets is pas toevallig als er ook toeval in zit. En twee of drie keer achter elkaar hetzelfde cijfer is best toevallig. Net als uitschieters in een normale verdeling.

Oefening

Schrijf een Pythonprogramma dat willekeurige getallen 1 t/m 9 genereert, uniform verdeeld. Bereken het product van die negen cijfers en bepaal het laatste (minst significante) cijfer. Herhaal dit duizend keer en houd bij hoe vaak elk cijfer als laatste voorkomt. Doe een voorspelling over hoe het histogram eruit komt te zien. Controleer of je voorspelling klopt.

49. Onafhankelijke fouten optellen

Als je een weegschaal hebt die ergens tussen en 0% en 2% te veel aangeeft, maar je weet niet hoeveel, wat moet je dan?

Over dat ‘ergens’ weet je niet zo veel. Domweg overnemen wat er op het display staat, zou niet slim zijn. Daar kun je voor corrigeren. Standaard 1% eraf trekken en een onzekerheid aanhouden van 1%, dat is geen slecht idee.

We wegen 406 gram suiker af. Dan zeggen we dat het 1% minder is en geven de bijbehorende onzekerheid op.

$$m_{\text{suiker}} = 0,402 \pm 0,004 \text{ kg}$$

Het zou kunnen, dat het inderdaad 406 gram was. Maar 398 gram is ook mogelijk.

Intussen hebben we een nieuwe weegschaal gekocht en de fabrikant is zo slim geweest om die ene procent al te corrigeren. Dat rekensommetje hoeft dus niet meer, maar die onzekerheid zit er nog steeds in. De schaal kan nu te veel of te weinig aangeven: we weten alleen niet hoeveel precies. (De fabrikant ook niet, anders hadden ze daar ook wel voor gecorrigeerd.)

We wegen opnieuw vier ons suiker af, maar ook nog eens twee ons bloem en een pond boter. (Toe maar!) Dat gooien we allemaal bij elkaar... hoeveel hebben we dan?

$$m_{\text{suiker}} = 0,404 \pm 0,004 \text{ kg}$$

$$m_{\text{bloem}} = 0,199 \pm 0,002 \text{ kg}$$

$$m_{\text{boter}} = 0,503 \pm 0,005 \text{ kg}$$

Het vervelende van die weegschaal is: we weten wel dat-ie afwijkt, maar niet precies hoeveel en ook niet of het te veel of te weinig is. (Hadden we nou die oude weegschaal nog maar! Die week meer af, maar je wist tenminste zeker dat je nooit te veel nam.)

Die schaal geeft dus structureel te veel of te weinig aan en we ontkomen er niet aan om de onzekerheden bij elkaar op te tellen. Het goede nieuws is: relatief blijft de onzekerheid gewoon die ene procent.

$$m_{\text{alles}} = 1,106 \pm 0,011 \text{ kg}$$

Als we twee keer zo veel hadden genomen, hadden we een significant cijfer moeten inleveren. Eigenlijk is dat niet logisch, maar dat zijn de regels.

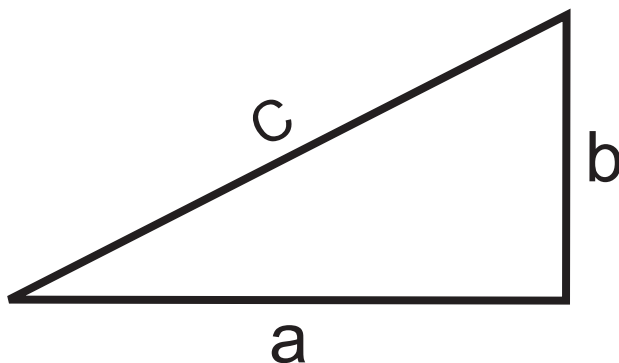
$$m_{dubbel} = 2,21 \pm 0,02 \text{ kg}$$

Nu gaan we iets anders doen. We lenen de weegschaal van onze burens. Twee verschillende weegschalen, van hetzelfde merk en met dezelfde specificaties. Die geven beide een procent te veel of te weinig aan, maar het probleem is: het kan best zijn dat de ene altijd te veel aangeeft en de andere altijd te weinig. Of andersom. Of beide te veel, of beide te weinig...

We gaan nu de suiker en de bloem in de ene weegschaal wegen en de boter in de andere. Hoe zit dat nu met het totaal? We moesten eerst de onzekerheden domweg optellen omdat het als maar erger werd. Maar het zou nu kunnen zijn dat de ene te veel aangeeft en de andere te weinig. Het is wel erg pessimistisch om ervan uit te gaan dat je het moet optellen.

Wiskundig gezien kun je bewijzen dat het *kwadratisch optellen* van de onzekerheden een betere maat geeft dan 'gewoon' optellen. Dat heeft iets met Taylorreeksen en orthogonaal zijn te maken en we gaan dat hier niet bewijzen. Hughes & Hase (2010) doen dat wel, dus als je wiskundig onderlegd bent...

Het bewijs hoef je niet te kennen of te begrijpen. Maar het is wel handig om in te zien dat een kwadratische optelling altijd minder oplevert dan een gewone optelling (tenzij een van beide nul is, maar dan tel je niets op).



Alleen als b of c gelijk is aan 0 is a gelijk aan de som van die twee. In alle andere gevallen is de waarde kleiner. Een rechte lijn tussen twee punten is nu eenmaal altijd een kortere weg dan via een omweg, met een haakse bocht erin.

50. Met én zonder afhankelijkheid

De bewegingsvergelijking van een vallende steen luidt als volgt.

$$s = \frac{1}{2}gt^2$$

We kiezen omlaag als positief en het punt van loslaten als oorsprong. De massa van de steen is zo groot dat de wrijvingskracht te verwaarlozen is ten opzichte van de zwaartekracht op de steen.

We willen de versnelling van de zwaartekracht berekenen en herschrijven daarom de formule in onderstaande vorm.

$$g = 2 \frac{s}{t^2}$$

Hoe moeten we nu de onzekerheid bepalen? Is er hier sprake van afhankelijke of onafhankelijke onzekerheden? De tijd wordt gemeten met een stopwatch of een camera en de afstand met een rolmaat of laserafstandmeter. Die metingen hebben niet veel met elkaar te maken. Je zou dus kunnen zeggen dat de fouten onafhankelijk zijn. De formule zou dan als volgt luiden.

$$\frac{\delta g}{|g|} = \sqrt{\left(\frac{\delta s}{s}\right)^2 + 2\left(\frac{\delta t}{t}\right)^2}$$

De factor 2 kent geen onzekerheid, dus zie je in de formule niet terug.

Als je pessimistischer bent en met afhankelijke onzekerheden wilt rekenen, kom je niet met een kwadratische optelling, maar een “gewone”.

$$\frac{\delta g}{|g|} = \frac{\delta s}{|s|} + 2 \frac{\delta t}{|t|}$$

Als de onzekerheid in de afstand 1% is en de onzekerheid in de tijd is 2%, dan zou uit deze formule volgen dat de relatieve onzekerheid in g gelijk is aan $1\% + 2 \times 2\% = 5\%$. Tel je die termen kwadratisch op, dan kom je op andere waarden.

$$\begin{aligned}\frac{\delta g}{|g|} &= \sqrt{(0,01)^2 + 2 \cdot (0,02)^2} = \sqrt{0,0001 + 2 \cdot 0,0004} \\ &= \sqrt{0,0009} = 3\%\end{aligned}$$

Welke berekening is correct? Het antwoord is: geen van beide.

De fout in s is onafhankelijk van de fout in t , maar de fout in t is *volledig* afhankelijk van zichzelf.

Een beetje filosofische kwestie misschien... Kan iets afhankelijk zijn van zichzelf? Als we zeggen dat x afhankelijk is van y , dan bedoelen we dat een verandering in x 'automatisch' een verandering in y betekent. Als x en y beide gelijk zijn aan t , dan is dat meeveranderen wel erg sterk aanwezig.

Hoe lossen we dit op? Het beste is om even uit te gaan van een andere grootte, die gelijk is aan het kwadraat van de tijd.

$$T = t^2$$

Daarvoor geldt:

$$\frac{\delta T}{|T|} = 2 \frac{\delta t}{|t|}$$

Als we dit invullen in de oorspronkelijke formule, dan staat er:

$$g = 2 \frac{s}{T}$$

Voor de onzekerheid geldt dan

$$\begin{aligned} \frac{\delta g}{|g|} &= \sqrt{\left(\frac{\delta s}{s}\right)^2 + \left(\frac{\delta T}{T}\right)^2} = \sqrt{\left(\frac{\delta s}{s}\right)^2 + \left(2 \frac{\delta t}{t}\right)^2} = \sqrt{\left(\frac{\delta s}{s}\right)^2 + 4 \left(\frac{\delta t}{t}\right)^2} \\ &= \sqrt{(0,01)^2 + 4(0,02)^2} = \sqrt{17}\% \approx 4,1\% \end{aligned}$$

De waarheid ligt dus in het midden. Nou ja, niet precies in het midden, maar wel tussen die 3% en 5%. Helemaal afhankelijk is wat te pessimistisch, en helemaal onafhankelijk is te optimistisch.

51. (On)afhankelijke normale verdelingen

Dat van die afhankelijkheid van verdelingen zou je zomaar voor waar kunnen aannemen, maar er is wiskundig iets eenvoudigs over te zeggen dat het inzichtelijker maakt.

Hughes & Hase (2010) vermelden op hun formuleblad (binnenzijde van de achterkant van het boek) dat de onzekerheid in de som van twee grootheden gelijk is aan de som van die twee onzekerheden, *kwadratisch opgeteld*.

Taylor (1997) vermeldt op zijn formuleblad (binnenzijde van de voorkant van het boek) hetzelfde, maar zegt daarbij dat dit alleen geldt voor *independent random errors*: onafhankelijke toevallige fouten. Wat altijd geldt, is dat die onzekerheid maximaal gelijk is aan de “gewone” som van de onzekerheden.

De boeken van Taylor en Hughes & Hase worden in *Metten & Weten* als referentie gebruikt, omdat ze beide als literatuur hebben gediend voor het vak Error Budgeting. Ze dienen dus regelmatig als toetssteen om te kijken of het wel klopt wat er in *Metten & Weten* staat. Het boek van Berendsen, oorspronkelijk in het Nederlands uitgegeven onder de (leukere) titel *Goed meten met fouten*, bevat soms aardige bewijzen.

Berendsen (2011) geeft in bijlage A1 met als titel *Combining uncertainties* een wiskundige analyse. De vetgedrukte vraag luidt: *Why do squared uncertainties add up in sums?*

Klopt dit wel? ‘Waarom tellen fouten kwadratisch op?’ luidde het Nederlandse origineel in zijn boek uit 1997. Dat klinkt al logischer. Al moet je natuurlijk zeggen dat de fouten zelf niet optellen: ze *worden* opgeteld. De eigenlijke vraag is: moet je onzekerheden kwadratisch optellen? En waarom dan, en wanneer? Daar gaat dit hoofdstuk over.

Als je twee normaal verdeelde grootheden hebt, geldt voor de verwachtingswaarden en de spreidingen:

$$\begin{aligned} E[x] &= \mu_x & E[(x - \mu_x)^2] &= \sigma_x^2 \\ E[y] &= \mu_y & E[(y - \mu_y)^2] &= \sigma_y^2 \end{aligned}$$

Je kunt die twee normaal verdeelde grootheden optellen.

$$z = x + y$$

Voor het gemiddelde van de verdeling van die z geldt dan:

$$\mu_z = \mu_x + \mu_y.$$

Nu de variantie. Die is gedefinieerd als het gemiddelde van het kwadraat van de afwijking van het gemiddelde. Dat is te schrijven als een verwachtingswaarde.

$$\sigma_z^2 = E[(z - \mu_z)^2]$$

Die z en μ_z kunnen we vervolgens vervangen door de formules hierboven.

$$\sigma_z^2 = E[(x + y - (\mu_x + \mu_y))^2]$$

Verander nu de volgorde van de termen tussen de ronde haken (en let op de plus- en mintekens).

$$\sigma_z^2 = E[(x - \mu_x + y - \mu_y)^2]$$

Nu is er een handige rekenregel voor kwadrateren van een optelling.

$$(A + B)^2 = A^2 + B^2 + 2AB$$

Pas die toe op bovenstaande, en je krijgt de uitkomst.

$$\sigma_z^2 = E[(x - \mu_x)^2 + (y - \mu_y)^2 + 2(x - \mu_x)(y - \mu_y)]$$

Deze uitkomst bevat twee ‘delen’: de eerste twee termen zijn de som van de beide varianties, de laatste term is het dubbele van de covariantie van beide grootheden. (En dan vermenigvuldigd met een factor 2.)

$$\sigma_z^2 = \sigma_x^2 + \sigma_y^2 + 2cov(x, y)$$

De variantie van de som van x en y is dus afhankelijk van twee termen die we kennen als de normale verdelingen beschreven zijn. Een normale verdeling ligt immers vast als je de μ en σ kent.

Hoe groot kan die derde term zijn, de covariantie?

$$cov(x, y) = E[(x - \mu_x)(y - \mu_y)]$$

Wat is de verwachtingswaarde van het product van de afwijkingen van het gemiddelde van beide grootheden?

Als de waarden van x en y op zuiver toeval berusten, dan zijn hun afwijkingen van het gemiddelde volstrekt ongecorreleerd. Gemiddeld zijn de afwijkingen gelijk aan nul, soms zijn ze positief en soms zijn ze negatief. Soms

zijn ze groot en soms zijn ze klein. Er zit geen relatie tussen en gemiddeld is hun product gelijk aan nul.

Maar wat nu als de waarden van x en y van elkaar afhankelijk zijn? De grootst denkbare afhankelijkheid is dat ze altijd gelijk zijn aan elkaar. In dat geval zijn ook hun gemiddelden en standaardafwijkingen aan elkaar gelijk en mag je voor y ook x invullen en μ_x voor μ_y .

$$\text{cov}(x, y) = E[(x - \mu_x)(x - \mu_x)] = E[(x - \mu_x)^2] = \sigma_x^2$$

Dan geldt dus dat

$$\sigma_z^2 = \sigma_x^2 + \sigma_y^2 + 2\text{cov}(x, y) = \sigma_x^2 + \sigma_x^2 + 2\sigma_x^2 = 4\sigma_x^2$$

Wat overblijft, is heel eenvoudig: $\sigma_z = \sqrt{4\sigma_x^2} = 2\sigma_x = \sigma_x + \sigma_y$.

We hebben nu bewezen dat de variantie van z

- minimaal gelijk is aan de kwadratische optelling en
- maximaal gelijk aan de gewone optelling

van de varianties van x en y .

Bij volstrekte onafhankelijkheid geldt het eerste.

Bij volstrekte afhankelijkheid geldt het eerste.

Bij gedeeltelijke afhankelijkheid zit het ertussenin.

Begrippen

52. Onzekerheid

Misschien wel het belangrijkste begrip in dit hele boek is onzekerheid.

Helemaal zeker is dat niet, maar *waarschijnlijk* is onzekerheid het belangrijkste begrip. Hoe weet je eigenlijk of een begrip het belangrijkste is? Kun je dat uitdrukken in een getal? Nee. Je zou dus kunnen zeggen dat een begrip in een boek het belangrijkste is als je geen ander begrip kunt bedenken dat belangrijker is. Maar wat nu als in een commissie van wijze mensen (zes personen) er vijf zitten die zonder enige twijfel stellen dat *onzekerheid* het belangrijkste is, en eentje beweert glashard dat *weten* belangrijker is. Ze weten het alle zes zeker, zeggen ze...

Wat is eigenlijk zekerheid? Heb je dat nog nooit meegemaakt dat je iets zeker meende te weten wat later toch niet waar bleek te zijn?
“Ik zou toch zweren dat...”

De GUM besteedt heel paragraaf 2.2 hieraan: *The term “uncertainty”*.

The word “uncertainty” means doubt, and thus in its broadest sense “uncertainty of measurement” means doubt about the validity of the result of a measurement. Because of the lack of different words for this general concept of uncertainty and the specific quantities that provide quantitative measures of the concept, for example, the standard deviation, it is necessary to use the word “uncertainty” in these two different senses.

We gaan hier vaak aan voorbij, maar ze hebben wel gelijk. Je pakt een meetlat van 100 cm en stelt vast dat de lengte van een ijzeren staaf 45,2 cm is. De onzekerheid schat je op 1 mm. Maar als je naar huis fietst en je meetresultaat wilt meenemen in een berekening, ga je je opeens afvragen of je de meetlat niet achterstevoren had. Dan was die staaf dus $100 - 45,2 = 54,8$ cm.

Er is nu *onzekerheid over de meting*. Misschien is die niet goed uitgevoerd en klopt het voor geen meter. Een andere mogelijkheid is dat je je blind hebt zitten staren naar die twee millimeterstreepjes en niet in de gaten had dat het geen 45,2 maar 46,2 was. Het kan ook zo zijn dat je afgeleid was toen je het opschreef en 45,2 noteerde terwijl je 54,2 had afgelezen. Maar ook als je het wel zeker weet (wanneer is dat zo) blijft er een stuk onzekerheid over. Dat is het soort onzekerheid dat we meestal bedoelden.

Twee soorten onzekerheid dus.

1. Of het wel klopt.
2. Hoeveel je er naast zou kunnen zitten.

The formal definition of the term “uncertainty of measurement” developed for use in this Guide and in the VIM [6] (VIM:1993, definition 3.9) is as follows:

uncertainty (of measurement)

parameter, associated with the result of a measurement, that characterizes the dispersion of the values that could reasonably be attributed to the measurand

Wat is de vertaling daarvan volgens Google Translate?

ölçülen büyüklüğe makul bir şekilde atfedilebilecek değerlerin dağılımını karakterize eden, bir ölçümün sonucuyla ilişkili parametre.

Die vertaling çlipt als een büş, maar omdat dit boek in het Nederlands geschreven is, kunnen we beter een andere doeltaal kiezen.

parameter, geassocieerd met het resultaat van een meting, die de spreiding karakteriseert van de waarden die redelijkerwijs aan de meetgrootte kunnen worden toegeschreven

De uncertainty is dus een getal (meestal met een eenheid), maar kan ook een percentage of fractie zijn.

Nu moeten we niet in de valkuil van dat woord ‘spreiding’ (*dispersion*) trappen. Hier staat NIET dat er DUS een spreiding in de METINGEN is. Het kan best zijn dat je de meting tien keer over doet en tien keer exact hetzelfde vindt. In het geval van onze ijzeren staaf: 45,2 cm. Maar er is wel een spreiding in de waarden die *redelijkerwijs de werkelijke waarde kunnen zijn*. Misschien is die staaf wel 45,194720 cm $\pm$ 5 μ m. Dat blijkt misschien wel als we een meetapparaat hebben in micrometers. Maar er zijn oneindig veel meer waarden die redelijk zijn: 45,226 cm kan ook. Of 45,24 cm. Er zit een spreiding in de mogelijkheden. Niet per se in de meting als we die herhalen.

The parameter may be, for example, a standard deviation (or a given multiple of it), or the half-width of an interval having a stated level of confidence.

Dit is even wennen als je je voor het eerst in dit vakgebied verdiept. De parameter kan een standaardafwijking zijn, maar ook een op het oog harde grens. Als je zeker denkt te weten dat een waarde tussen twee streepjes zit (tussen 45,1 en 45,3 cm), dan kan het niet zo zijn dat de werkelijke waarde 45,4 cm is. Maar stel nu eens dat die lengte uit een grote reeks normaal verdeelde metingen kwam. Dan zou die 0,1 cm de standaardafwijking zijn. Het is niet onmogelijk dat de werkelijke waarde dan 45,4 cm was. De kans dat je twee keer de standaardafwijking van het gemiddelde zit, is ongeveer 5%. Niet heel groot, maar wel iets wat gemiddeld één op de twintig keer voorkomt.

Laten we nog even verder lezen in de GUM.

The definition of uncertainty of measurement given in 2.2.3 is an operational one that focuses on the measurement result and its evaluated uncertainty. However, it is not inconsistent with other concepts of uncertainty of measurement, such as

- *a measure of the possible error in the estimated value of the measurand as provided by the result of a measurement;*
- *an estimate characterizing the range of values within which the true value of a measurand lies (VIM:1984, definition 3.09).*

Although these two traditional concepts are valid as ideals, they focus on unknowable quantities: the “error” of the result of a measurement and the “true value” of the measurand (in contrast to its estimated value), respectively. Nevertheless, whichever concept of uncertainty is adopted, an uncertainty component is always evaluated using the same data and related information.

Volgens Google Translate:

De definitie van meetonzekerheid gegeven in 2.2.3 is een operationele definitie die zich richt op het meetresultaat en de geëvalueerde onzekerheid ervan. Het is echter niet inconsistent met andere concepten van meetonzekerheid, zoals

- *een maatstaf voor de mogelijke fout in de geschatte waarde van de maatstaf, zoals blijkt uit het resultaat van een meting;*
- *een schatting die het bereik van waarden karakteriseert waarbinnen de werkelijke waarde van een meetgrootte ligt (VIM:1984, definitie 3.09).*

Hoewel deze twee traditionele concepten geldig zijn als idealen, richten ze zich op onkenbare grootheden: respectievelijk de ‘fout’ van het resultaat van een meting en de ‘echte waarde’ van de meetgrootte (in tegenstelling tot de geschatte waarde ervan). Niettemin wordt, welk concept van onzekerheid ook wordt gehanteerd, een onzekerheidscomponent altijd geëvalueerd met behulp van dezelfde gegevens en gerelateerde informatie.

Op die onkenbare grootheden, namelijk ‘werkelijke waarde’ en ‘fout’ komen we later nog terug.

53. Verificatie & Validatie

Het verschil tussen *verificatie* en *validatie* zou je kunnen opzoeken in een woordenboek, maar dan kom je misschien niet veel verder.

- *Verifiëren* is “de juistheid (of waarheid) aantonen”.
- *Valideren* is “geldig verklaren”.

Beide woorden komen uit het Frans, vijftiende en zestiende eeuw. Maar ja, wat heb je eraan om dat te weten?

Wikipedia geeft fraai weer wat het verschil is tussen *verificatie* en *validatie*.

- *Verificatie: men controleert of het systeem juist is gemaakt (zijn alle eisen verwerkt?)*
- *Validatie: men controleert of het juiste systeem is gemaakt (doet het wat de klant wil?)*

Dat kun je nog korter zeggen:

- *Verificatie: maken we het juist? (klopt het met de eisen?)*
- *Validatie: maken we het juiste? (kloppen de eisen?)*

Ik kan de breedte van een rechthoekige tafel opmeten met een rolmaat. Die meting kan ik laten verifiëren door iemand anders te vragen om het met een duimstok na te meten. Vindt hij of zij hetzelfde? Dan is er sprake van *verificatie* van de meting. Als die ander nagaat of ik wel de kleinste van beide zijden heb genomen, is er sprake van *validatie*. De langste zijde moet je namelijk niet hebben: dat is niet breedte, maar de lengte.

Er is een verschil tussen *nagaan of* en *nagaan dat*. In het eerste geval kijk je *of iets zo is* (of niet!), in het tweede geval zie je *dat het zo is*.

“Heb je gecontroleerd of de deur op slot zit?”

“Ja.”

“En? Zat-ie op slot?”

“Nee.”

“Moet je hem niet op slot doen dan?”

“Nee.”

“Waarom niet?”

“Omdat ik dat gedaan heb, toen ik zag dat-ie niet op slot zat...”

“Toen”? Bedoel je niet: ‘nadat’?”

54. Relatief & Absoluut

De *absolute onzekerheid* is een grootte die wordt uitgedrukt in dezelfde eenheden als de gemeten grootte. Die geeft dus het bereik aan waarin de werkelijke waarde waarschijnlijk zal liggen. De werkelijke waarde kan zich buiten dat bereik bevinden.

Let op dat het bereik waarover we het nu hebben dus *twee keer zo groot* is als de grootte die we onzekerheid noemen. En we weten nog niet eens voor 100% zeker dat de grootte daarbinnen ligt.

De *relatieve onzekerheid* is de absolute onzekerheid gedeeld door de gemeten grootte zelf. Dat is dus een dimensieloos getal, vaak aangeduid als een percentage.

Let op dat dit percentage groter kan zijn dan 100%.
Een standaardvoorbeeld is een spanning die heel klein is. Die spanning zou bijvoorbeeld -2 mV kunnen zijn terwijl de onzekerheid 5 mV is. Dan is de relatieve onzekerheid 40%.
(Ook een relatieve onzekerheid is een absolute waarde, nooit negatief.)

Kan een relatieve fout ook *relatief* groot of klein zijn? Ja. In hoofdstuk 50 (titel: *Met én zonder afhankelijkheid*) wordt gerekend aan de onzekerheid in de gravitatieversnelling die je vindt via een valproef. Gekeken wordt of ze afhankelijk of onafhankelijk van elkaar zijn en hoe dat dan doorwerkt. Je ziet daar ook dat de onzekerheid volgt uit een formule waarin de relatieve fout in de tijd twee keer meetelt en de relatieve fout in de afstand één keer.

Laten we voor het gemak even doen alsof die onzekerheid in de lengte en in de tijd afhankelijk van elkaar zijn.

Waarschijnlijk zijn ze dat *niet*. Die lengte meet je met een rolmaat en de tijd met een stopwatch. Het is niet logisch dat de onzekerheid van de ene iets met de andere te maken. Dat de slingertijd groter wordt als je de slinger langer maakt, is wel waar. Maar we hebben het niet over de afhankelijkheid in de werkelijke waarden, maar over de afhankelijkheid in de onzekerheden.

Dan is dit de formule.

$$\frac{\delta g}{|g|} = \frac{\delta s}{|s|} + 2 \frac{\delta t}{|t|}$$

Kun je zeggen welke onzekerheid meer bijdraagt in de totale onzekerheid? Nee. Je weet wel dat het *gewicht* bij de tijd twee keer zo groot is. Maar als de

relatieve onzekerheid in de afstand 4% is en de relatieve onzekerheid in de tijd 0,1%, dan draagt die eerste een factor twintig keer zo veel bij.

Kun je zeggen dat de ene relatieve fout groter is dan de andere? Ja. Die 4% is veertig keer zo veel als 0,1%.

Dat is het aardige aan relatieve onzekerheden. Ze geven niet alleen een indruk van hoe de onzekerheid zich verhoudt tot de waarde, maar je kunt ze ook vergelijken met andere grootheden met een andere dimensie. Door het relatieve ben je van die dimensies af.

Relatieve onzekerheden hebben ook iets vervelends. Voor kleine waarden van de gemeten grootheden worden ze al gauw erg groot. Als ik wil vaststellen dat de spanning ‘bijna nul’ is, zeg ik liever dat de onzekerheid 2 mV is dan dat de onzekerheid 40% is. Dat zegt namelijk weinig.

Kun je zeggen dat de ene relatieve fout relatief groot is? Ja. Hij is bijvoorbeeld veertig keer zo groot als die andere. Nu kun je óók nog van mening verschillen over de vraag hoeveel keer zo groot relatief groot is. Alles is betrekkelijk.

55. Werkelijke waarde

De makers van de GUM zijn wel grappenmakers! In hun document dat de basis legt voor het vakgebied rondom onzekerheden lezen we op pagina 16:

NOTE The term “true value” (see Annex D) is not used in this Guide for the reasons given in D.3.5; the terms “value of a measurand” (or of a quantity) and “true value of a measurand” (or of a quantity) are viewed as equivalent.

Het is zoiets als ‘zeg nooit “nooit”’. Ze noemen datgene wat ze niet gebruiken al twee keer in deze ene alinea. Maar het wordt pas echt lollig als je telt hoe vaak de term *true value* voorkomt in hun hele *manual*: 55 keer. Dat is toch niet heel weinig in een document van 134 pagina’s. In ieder geval vaker dan nooit.

Nu gebruiken ze de term meestal om te zeggen dat ze hem niet gebruiken. Als dat niet onder ‘gebruiken’ valt, valt er wat voor te zeggen dat je het toch opschrijft.

Het woord *value* is vrij eenvoudig en eenduidig in het Nederlands te vertalen, namelijk met ‘waarde’. Bij *true* ligt dat iets lastiger. Het is ‘waar’, ‘juist’, ‘echt’, ‘werkelijk’, ‘kloppend’, ‘correct’, maar ook ‘trouw’, ‘eerlijk’, ‘zuiver’, ‘oprecht’, ‘waarachtig’ en ‘waarheidsgetrouw’. Je zou *true value* kunnen vertalen met ‘echte waarde’, maar de voorkeur gaat uit naar ‘werkelijke waarde’.

Eigenlijk is ‘echte waarde’ de vertaling van *real value*.

Je zou ook ‘ware waarde’ kunnen zeggen in plaats van ‘werkelijke waarde’, maar dat klinkt raar.

Het is trouwens niet echt waar dat de makers van de GUM grappenmakers zijn. Het lezen van hun document zet je maar zelden aan het lachen. Waar ze wel goed in zijn, is sommige dingen heel zorgvuldig onder woorden brengen. Kijk maar naar de cursieve *NOTE* hierboven. Ze hebben volkomen gelijk dat ‘de werkelijke waarde’ niet meer is dan ‘de waarde’. Het is gewoon de waarde die je probeert te meten. Uiteraard is dat de werkelijke waarde... Je zou wel gek zijn als je de *onwerkelijke* waarde probeerde te meten.

“Even kijken hoe lang dit koord *niet* is.”

“En?”

“Even kijken... Ik lees af dat het bijna 60 cm is. Dus laten we zeggen: twee meter.”

“Twee meter?”

“Ja. Ik wou namelijk weten hoe lang het *niet* is.”
 “O ja, dat zei jij al. Wil je het echt niet weten?”
 “Ja.”
 “Is dat waar?”
 “Ja.”
 “Is dat écht waar?”
 “Ja. Iets kan niet waar zijn en tegelijkertijd niet écht waar...”
 “Da’s waar. En dan snap ik waarom je twee meter zegt. Als je 62 cm gezegd had, liep je het gevaar dat het toch nog waar was.”
 “Ik meende het niet hoor. Het was maar een grap...”
 “Echt?”
 “Ik heb het tegendeel nooit ontkend.”

In dit boek gebruiken we vaak de term ‘werkelijke waarde’. Dat is de waarde die de te meten grootte heeft. Het woord ‘werkelijk’ gebruiken we om het te onderscheiden van de *gemeten waarde* en de *geaccepteerde waarde*.

Wat van belang is, is dat je weet met welke waarde je te maken hebt.

- *De gemeten waarde.* Als je de rustmassa van een elektron zelf thuis op een zolderkamer probeert te bepalen, kom je bijvoorbeeld op een meetresultaat van $(10,2 \pm 1,5) \cdot 10^{-31}$ kg.
- *De geaccepteerde waarde.* De rustmassa van het elektron bedraagt volgens ‘de boeken’ (het *National Institute of Standards and Technology*) $9,10938356(11) \times 10^{-31}$ kg. Dat is niet strijdig met jouw eigen meting. De onzekerheid hierin is wel veel kleiner.
- *De werkelijke waarde.* De rustmassa van het elektron is in werkelijkheid... onbekend. We mogen aannemen dat het ligt in het gebied dat rondom de geaccepteerde waarde is opgeven. Maar als je beweert dat die rustmassa exact gelijk is aan $9,10938356 \times 10^{-31}$ kg, dan heb je ongelijk.

De werkelijke waarde weet niemand. Althans: geen mens op aarde. Als je in God gelooft, in een Schepper van alles, kun je nadenken over de vraag of Hij het weet. Van God wordt beweerd dat Hij niet alles kan. Hij kan bijvoorbeeld niet liegen. Kan Hij wel alles weten? Ik weet het niet. Ik ben maar een mens. Ik ken de werkelijke waarde van de rustmassa van het elektron al niet eens. Terwijl dat best een klein getal is.

56. Fout

error (of measurement)

result of a measurement minus a true value of the measurand

Als je gelezen hebt dat de makers van de GUM beweren dat ze nooit de term *true value* gebruiken, sta je van hun definitie van *error* wel te kijken. Vooral omdat ze het in de eerstvolgende noot ook weer doen.

NOTE 1 Since a true value cannot be determined, in practice a conventional true value is used.

NOTE 2 When it is necessary to distinguish “error” from “relative error”, the former is sometimes called absolute error of measurement. This should not be confused with absolute value of error, which is the modulus of the error.

Als de fout (*error*) het verschil is tussen de gemeten waarde en de werkelijke waarde, ligt het voor de hand dat ook voor de fout geldt dat je die niet kent. De fout is: hoeveel je afwijkt van de werkelijke waarde. Je meetresultaat ken je exact, want je schrijft het in een eindig aantal decimalen op. Maar die *true value*... kenden we die maar.

In hoofdstuk 3.2, Errors, effects, and corrections, is de GUM (terecht!) zorgvuldig in het gebruik van het woord ‘error’.

3.2.1 In general, a measurement has imperfections that give rise to an error (B.2.19) in the measurement result. Traditionally, an error is viewed as having two components, namely, a random (B.2.21) component and a systematic (B.2.22) component.

NOTE Error is an idealized concept and errors cannot be known exactly.

Zucht. Ik vind het zo vermoeiend lezen als er dingen als (B.2.19) en (B.2.22) in een zin staan. Hadden ze daar nou geen voetnoten voor kunnen gebruiken? Ja, dat had gekund. Maar ze hebben het niet gedaan. In *Metten & Weten* staan trouwens nooit voetnoten.¹

Het zijn dus toevallige en systematische fouten die zorgen voor een onzekerheid. Het verschil tussen de werkelijke waarde en je meting kan per

¹ Behalve dan deze ene.

geval verschillen, maar ook de hele tijd hetzelfde zijn.
Je kent het verschil niet. Vandaar die onzekerheid.

3.2.2 Random error presumably arises from unpredictable or stochastic temporal and spatial variations of influence quantities. The effects of such variations, hereafter termed random effects, give rise to variations in repeated observations of the measurand. Although it is not possible to compensate for the random error of a measurement result, it can usually be reduced by increasing the number of observations; its expectation or expected value (C.2.9, C.3.1) is zero.

Je kunt niet systematisch compenseren voor een toevallige fout, want die is telkens anders. Maar je kunt óók niet compenseren voor een systematische fout als je die niet kent.

NOTE 1 The experimental standard deviation of the arithmetic mean or average of a series of observations (see 4.2.3) is not the random error of the mean, although it is so designated in some publications. It is instead a measure of the uncertainty of the mean due to random effects. The exact value of the error in the mean arising from these effects cannot be known.

Ook dit is waar. Je kunt de standaardafwijking van de theoretische verdeling natuurlijk schatten uit herhaalde metingen, maar je weet niet exact hoe groot die in werkelijkheid is. Als je heel veel metingen doet, kom je zeer waarschijnlijk wel heel dicht in de buurt. Er zit niet alleen een fout in je beste schatting, maar ook in je onzekerheid.

NOTE 2 In this Guide, great care is taken to distinguish between the terms “error” and “uncertainty”. They are not synonyms, but represent completely different concepts; they should not be confused with one another or misused.

Het staat er met nadruk: *great care is taken*. Het is echt iets anders, ook al heeft het iets met elkaar te maken. Het zijn *completely different concepts*. Toch kom je ze vaak tegen: mensen die bijvoorbeeld zeggen dat de *fout* in een meting 2 mm is. Als je dat weet, waarom compenseer je er dan niet voor? Antwoord: omdat het niet de *fout* is, maar de *onzekerheid* die ze ten onrechte fout noemen.

NOTE The uncertainty of a correction applied to a measurement result to compensate for a systematic effect is not the systematic error, often termed bias, in the measurement result due to the effect as it is sometimes called. It is instead a measure of the uncertainty of the result due to incomplete knowledge of the required value of the correction. The error arising from imperfect compensation of a systematic effect cannot be exactly known. The terms “error” and “uncertainty” should be used properly and care taken to distinguish between them.

57. Type A & Type B

De GUM (*Guide to the Expression of Uncertainty in Measurement*) zou je kunnen beschouwen als de standaard rondom het noteren van onzekerheden. In dat document legt de ISO een aantal zaken vast, waaronder ook terminologie.

Ratcliffe & Ratcliffe (2014) beginnen hun boek met een hoofdstuk over terminologie en besteden daarin aandacht aan het feit dat de GUM over Type A en Type B spreekt. Ze merken daarbij op dat de GUM zelf óók gebruik maakt van de als verouderd beschouwde termen.

De tekst uit de GUM (ISO, 2008) waar het om gaat, is als volgt.

The uncertainty in the result of a measurement generally consists of several components which may be grouped into two categories according to the way in which their numerical value is estimated:

A. those which are evaluated by statistical methods,

B. those which are evaluated by other means.

There is not always a simple correspondence between the classification into categories A or B and the previously used classification into “random” and “systematic” uncertainties. The term “systematic uncertainty” can be misleading and should be avoided.

Any detailed report of the uncertainty should consist of a complete list of the components, specifying for each the method used to obtain its numerical value.

Kirkup & Frenkel (2006) geven een aardig voorbeeld om dit subtiele verschil duidelijker te maken, aan de hand van een meetinstrument waarbij er sprake is van een *drift*. In het Nederlands wordt *drift* ook wel *verloop* genoemd. Het is een kleine continue verandering in de weergegeven meetwaarden in een bepaald tijdsverloop waarbij de te meten waarden constant blijven. Dat kan een gevolg zijn van een verloop in de temperatuur of het inzakken van de spanning van een batterij.

Als er sprake is van ruis in het meetinstrument en de drift is betrekkelijk klein, dan is die niet zomaar waar te nemen. Een weegschaal die per minuut een gram meer aangeeft en geen toevallige fout heeft, is een voorbeeld daarvan. Als er geen ruis was, zou je de drift dus vrij eenvoudig kunnen vaststellen én ervoor corrigeren. Maar als de ruis groot is ten opzichte van de drift, wordt dat lastiger.

Hoe bepaal je een drift? Met *statistische methoden*. Als je elke seconde één meting zou kunnen doen, kon je de ruis uitmiddelen. Dat kan niet altijd,

bijvoorbeeld omdat de omstandigheden van de meting maken dat je niet vaak genoeg kunt herhalen. Er is dan sprake van een systematische fout en een toevallige fout. Beide zou je onder Type A moeten scharen.

De GUM geeft elders ook nog een toelichting op het verschil.

In some publications, uncertainty components are categorized as “random” and “systematic” and are associated with errors arising from random effects and known systematic effects, respectively. Such categorization of components of uncertainty can be ambiguous when generally applied. For example, a “random” component of uncertainty in one measurement may become a “systematic” component of uncertainty in another measurement in which the result of the first measurement is used as an input datum. Categorizing the methods of evaluating uncertainty components rather than the components themselves avoids such ambiguity. At the same time, it does not preclude collecting individual components that have been evaluated by the two different methods into designated groups to be used for a particular purpose.

Een toevallige waarde ligt vast vanaf het moment dat je die gemeten hebt. Als die toevallige vaste waarde vervolgens wordt meegenomen in een reeks berekeningen, levert dat een systematische fout op. De toevallige fout is daarmee systematisch geworden.

58. Meetschalen

In het wetenschappelijke tijdschrift *Science* van vrijdag 7 juni 1946 werd er een artikel gepubliceerd met de titel *On the Theory of Scales of Measurement*. Het was afkomstig van de psycholoog Stanley Smith Stevens van de universiteit van Harvard en beschreef de verschillende niveaus waarin je metingen kunt onderverdelen.

Hij benoemde vier verschillende meetschalen.

- *nominaal*: je kunt zeggen dat iets tot een bepaalde categorie behoort
- *ordinaal*: je kunt zeggen dat het ene meer of minder is dan het andere
- *interval*: je kunt zeggen hoeveel het ene meer of minder is dan het andere
- *ratio*: je kunt zeggen hoeveel het ene meer of minder is dan het andere en dat uitdrukken in een percentage

Noemenswaardig is ook nog de circulaire schaal. Als de ene klok op kwart over drie staat en de andere op kwart over negen, dan kun je niet vaststellen of de ene voorloopt op de andere of omgekeerd.

Een bijzonder geval van de nominale schaal is een binaire schaal. Je kunt bijvoorbeeld zeggen dat een getal even of oneven is, terwijl even niet meer of minder is dan oneven.

Voorbeelden van een ordinale schaal kom je in de dagelijkse praktijk vaak tegen. Als je kiest tussen aardbeienijs en chocolade-ijs, dan kun je een sterke voorkeur voor het ene of het andere hebben. Misschien ook wel of het veel of weinig uitmaakt. Maar daar hoort niet zomaar een getal bij. Of het zou een jury-cijfer moeten zijn.

Bij een intervalschaal hebben de getallen wel een zinvolle betekenis, maar drukken ze niet meer uit dan een verschil. Een korfbalwedstrijd die met 21-12 gewonnen werd, kende een ruimere overwinning dan 16-14. Maar kun je zeggen dat 21-12 beter is dan 16-6? Of bij een sport waar minder gescoord wordt: 3-1 is beter dan 2-1. Maar is 3-1 ook beter dan 1-0?

Bij een ratioschaal kun je alles wat er maar met getallen te bedenken is. Noa is 1,60 meter lang en Bo is 2 meter. Bo is dus langer. Het scheelt 40 centimeter. Bo is 25% langer dan Noa. Noa is 20% kleiner dan Bo. Dat percentage hangt af van wie je als uitgangspunt neemt, maar de getallen zijn hard.

In tabelvorm ziet dat er als volgt uit.

Niveau		Kenmerkend	Volgorde	Verschillen	Nulpunt	Operatoren
Categorisch, Kwaliteit	nominaal	x				=, ≠
	ordinaal	x	x			>, <
Kardinaal, Numeriek, Kwantitatief	interval	x	x	x		+, −
	ratio	x	x	x	x	×, /

Het verhaal ‘meetschaal’ is lastig te combineren met onzekerheden. Als ik twee metalen staven heb met lengtes 50 ± 6 cm en 48 ± 4 cm, dan kan niet eens met zekerheid vaststellen welke van de twee groter is. Over de *beste schatting* kun je zeggen dat het verschil 2 cm of 4 % is: dat is een ratioschaal.

Een interessant voorbeeld van meetschalen is temperatuur. Uitgedrukt in graden Celsius of Fahrenheit is dat een intervalschaal. In kelvin is het een ratioschaal, omdat het absolute nulpunt een echt nulpunt is.

Oefeningen

1. Vandaag staat de thermometer op $20,2^\circ\text{C}$ en gisteren gaf hij 20°C aan. Kun je zeggen dat het 1% warmer is dan gisteren?
2. Anne, Beau en Cor schaken tegen elkaar. Anne wint de partij tegen Beau. Beau wint daarna de partij tegen Cor. Cor wint tenslotte de laatste wedstrijd tegen Anne. Analyseer en beschrijf wat hier gebeurt.
3. Jan en Piet (uit Tilburg) stellen zich net als Kees (uit Eindhoven) verkiesbaar als voorzitter van Carnaval Brabant. Kees krijgt 37% van de stemmen, Piet 35% en Jan 27%. Omdat het verschil tussen Kees en Piet zo klein is, wordt er besloten om opnieuw te stemmen met alleen die twee kandidaten. Bij de tweede ronde krijgt Kees 41% van de stemmen en Piet 59%. Bedenk een verklaring hiervoor

59. Discrepantie

Taylor (1997) besteedt in zijn boek een aparte paragraaf aan het begrip *discrepancy*, in het Nederlands: discrepantie. En over Nederlands gesproken: laten we eens opzoeken wat het woordenboek van *Van Dale* erover zegt.

dis·cre·pan·tie

zelfstandig naamwoord • de v • discrepanties

1 onderlinge afwijking, het uiteenlopen

2 tegenstrijdigheid, tegenspraak

Twee verschillende betekenissen: een onderlinge afwijking is nog geen tegenstrijdigheid. Als iemand zegt dat iets nog ruim een kwartier duurt en een ander heeft het over bijna een half uur, dan zit er een afwijking tussen die onderlinge uitspraken. Maar er is nog geen sprake van een tegenstrijdigheid.

2.3 Discrepancy

Before I address the question of how to use uncertainties in experimental reports, a few important terms should be introduced and defined. First, if two measurements of the same quantity disagree, we say there is a discrepancy. Numerically, we define the discrepancy between two measurements as their difference.

discrepancy = difference between two measured values of the same quantity

More specifically, each of the two measurements consists of a best estimate and an uncertainty, and we define the discrepancy as the difference between the two best estimates.

Het is een beetje opletten met deze definitie. Want wat is het verschil tussen twee gemeten waarden? Strikt genomen is dat niet alleen het verschil tussen de waarden zelf, maar ook een berekening van de onzekerheid in dat verschil. Dat hebben we besproken in hoofdstuk 37. Als definitie houden we aan wat Taylor schrijft in de tekst achter ‘*More specifically...*’

Discrepantie is het verschil tussen twee beste schattingen van twee verschillende metingen aan dezelfde grootte.

Het voorbeeld dat Taylor noemt, geeft meer duidelijkheid.

For example, if two students measure the same resistance as follows

Student A: 15 ± 1 ohms

and

Student B: 25 ± 2 ohms,

their discrepancy is

$$\text{discrepancy} = 25 - 15 = 10 \text{ ohms.}$$

Recognize that a discrepancy may or may not be significant.

Is er sprake van een discrepantie tussen de metingen? Ja, en die discrepantie bedraagt 10 ohm. Maar stel nu eens dat er een nog een student een meting doet. Die noemen we Student C (niet genoemd in het boek van Taylor) en die vindt:

Student C: $(15,5 \pm 1,5) \Omega$

(Student C zet altijd haakjes om de getallen en voor de eenheid, en gebruikt het gestandaardiseerde symbool Ω in plaats van het woord *ohm*. Niet dat het fout is wat Taylor doet, maar wat Student C doet is wel netter.)

Nu de vraag: is er een discrepantie tussen de metingen van Student A en Student C. Het antwoord: ja, en die discrepantie bedraagt $0,5 \Omega$. Er is dus wél een *discrepantie* maar geen *onderlinge tegenstrijdigheid*. We gebruiken betekenis 1 uit *Van Dale*.

Hughes & Hase (2010) besteden ook aandacht aan *discrepantie*, maar hebben geen apart hoofdstuk met dat woord als titel en de term ontbreekt ook in de inhoudsopgave. Ook dat is niet fout, maar het vestigt minder de aandacht op de betekenis van het begrip.

3.3.4 Comparing experimental results with an accepted value

If we are comparing an experimentally determined value with an accepted value, we can use Table 3.1 to define the likelihood of our measurement being accurate. The procedure adopted involves measuring the discrepancy between the experimentally determined result and the accepted value, divided by the experimentally determined standard error. If your experimental result and the accepted value differ by:

- *up to one standard error, they are in excellent agreement;*
- *between one and two standard errors, they are in reasonable agreement;*
- *more than three standard errors, they are in disagreement.*

Hughes & Hase definiëren het begrip *discrepancy* dus niet, maar gebruiken het gewoonweg in de betekenis die Taylor hanteert en die wij in dit boek ook gebruiken.

Opmerkelijk is dat de auteurs een soort tussenvorm tussen strijdig en in overeenstemming hanteren. Beter gezegd: ze hebben gradaties in overeenstemming, namelijk *uitstekend* en *redelijk*.

Zijn ze dan niet veel te soepel? Zijn Hughes & Hase van die docenten die bij een 5,3 altijd wel een paar tienden zien te vinden? Nee, maar ze geven wel rekenschap van het feit dat metingen en onzekerheden meer met schattingen en kansen te maken hebben dan je wellicht denkt. Om een vergelijking te maken met die 5,3: een student die nét een onvoldoende haalt, beheerst de stof misschien toch wel goed genoeg. Daar staat tegenover dat je bij een student met 5,7 zo je vraagtekens kunt plaatsen. Eén vraag anders of één keer minder goed gegokt en het was onvoldoende geweest. Bij iemand die een 7,5 haalt, ben je veel zekerder van je zaak.

60. Relevant & Significant



Anne wil graag weten of het sierradendoosje dat ze kwijt is op de kast ligt. Dat kan ze niet zien, omdat ze niet lang genoeg is om zonder ergens op te gaan staan op de kast te kijken. Daarom vraagt ze het aan Wil, want die is langer. Helaas: Wil wil wel, maar kan het ook niet zien.

Is Wil dan niet langer dan Anne? Jawel. Maar het verschil is te klein om van belang te zijn in dit geval. Ze schelen ongeveer zes centimeter en ook als Wil nóg zes centimeter langer was, kon hij nog steeds niet op de kast kijken.

Het verschil is niet *relevant*: het maakt niet uit.

Of iets relevant is of niet, hangt meestal af van de context. Er zijn heus wel situaties waar het wél uitmaakt dat Wil langer is. Bijvoorbeeld als ze iets uit de kelder nodig hebben. Die is 1 meter 80 hoog en Anne kan daar staan zonder te hoeven bukken. Lekker makkelijk! En dus ook van belang.

Het *online* woordenboek van *Van Dale* kent de woorden *significant* en *relevant* beide.

re•le•vant (bijvoeglijk naamwoord)

1 van belang, van betekenis

maatschappelijk relevant: nuttig voor de samenleving

sig•ni•fi•cant (bijvoeglijk naamwoord, bijwoord)

1 statistisch niet aan toeval toe te schrijven en dus betekenisvol

De betekenissen bevatten allebei het woord 'betekenis' en zijn dus aan elkaar verwant. Kennelijk is 'van betekenis' niet hetzelfde als 'betekenisvol'...

Opvallend is dat *significant* een bijwoord kan zijn een *relevant* niet... Hoe zit dat? Een bijwoord kan betrekking hebben op een werkwoord: twee dingen kunnen *significant verschillen*, maar je zegt niet dat ze *relevant verschillen*.

Het woord 'significant' komt van het Latijnse *significare*, dat je kunt opdelen in *signum* (teken, denk aan het Engelse woord *sign*) en *ficare* (maken).

Het woord *relevant* komt ook uit het Latijn, maar dan van *relevare* dat 'verlichten' of 'helpen' betekent. En zo zijn we terug bij Anne en Wil, die wel wil helpen en als teken daarvan naar de kast toe loopt.

In de context van dit boek is het voorbeeld van Anne en Wil handig. We kunnen namelijk vrijwel met zekerheid zeggen dat Wil langer is dan Anne. Zes centimeter: het kan ook best vijf of vier zijn, maar heus wel meer dan één. Dat zie je meteen als ze naast elkaar gaan staan. Op haar naaldhakken is Anne net zo lang als Wil op z'n sloffen.

We *weten vrijwel zeker* dat Wil langer is, maar *het maakt niet uit* in de casus van de kast.

Oefening

Een student haalt op een tentamen een 5,4 en twee weken later op het her-tentamen een 5,5.

1. Leg uit waarom het verschil wel of niet *significant* is.
2. Leg uit waarom het verschil wel of niet *relevant* is.

61. Onzekerheid over de onzekerheid

Anne zegt dan wel dat ze 1 meter 77 is, maar ze overdrijft een beetje. Ze zou namelijk graag wat langer zijn. Even lang als Wil. Waarom ze dat wil? Tsja... Onzekerheid misschien. We weten het niet. Sommige mensen denken dat je méér bent met een grotere lengte, een hoger salaris of een lagere stem.

Wil gaat ervan uit dat er een aardige onzekerheid zit op de lengte die Anne en hij opgeven. Van hemzelf, omdat hij nooit serieus gemeten heeft. Van haar, omdat hij weet dat ze zich graag groter voordoet dan ze is. Daar zit bij beiden wel een marge van 3 cm op.

Als we het over de lengte van Anne hebben, moeten we dus twee getallen noemen: haar geschatte (of gemeten of opgegeven) lengte, in het Engels ook wel de *best estimate*, en de onzekerheid daarop.

De beste schatting is 1,77 m.

De onzekerheid is 0,03 m.

Anne zelf zegt dat de onzekerheid minder dan 1 cm is. Dat zie je vaak bij mensen die onzeker zijn: die zijn wat stelliger in hun uitspraken. Denk ik.

Er zit dus rek in die 1 meter 77, maar die is in ieder geval een stuk kleiner dan 10%. Niemand beweert dat Anne onder de 1 meter 64 zit en er zijn ook geen familieleden en kennissen die denken dat ze de 1 meter 94 haalt. Een paar vriendinnen noemen lengtes als 1 meter 80 of zelfs 1 meter 82, maar dat is nog steeds minder dan 5% van die inmiddels veelbesproken 1,77 m.

Hoe zit dat met de onzekerheid? Anne houdt het op 1 cm, Bill op 2 cm en Wil op 3 cm. Een altijd wat pessimistisch uitgevallen oom heeft ooit geroepen 'dat je daar zomaar twee duimen naast kunt zitten', maar ook die pakweg 5 cm maakt nog steeds niet dat ze langer is dan die ene vriendin beweerde.

We zagen dus al dat die lengte ruim binnen 10% te schatten is.

Maar hoe goed kennen we de onzekerheid? Als je die oom buiten beschouwing laat, zijn 1 t/m 3 cm getallen die voorbijkwamen. De onzekerheid zelf zou je dus kunnen noteren als (2 ± 1) cm.

De onzekerheid op de onzekerheid is minstens 50%.

Moraal van dit verhaal: het is veel moeilijk om te schatten hoe groot je onzekerheid is dan te schatten hoe groot iemand is.

62. Effectief & Efficiënt

Als je effe niet goed nadenkt, zou je de woorden *effectief* en *efficiënt* zomaar door elkaar kunnen halen. Dat is niet handig. Wat zijn ‘gewone’ Nederlandse synoniemen voor deze twee uit het Latijn geleende termen?

- *Effectief*: doeltreffend
- *Efficiënt*: doelmatig

Dat zijn twee synoniemen, maar het is nog geen uitleg.

- Iets is *doeltreffend* als je het doel treft: je bereikt wat je wilt hebben.
- Iets is *doelmatig* als je dat doel met weinig middelen of moeite bereikt.

Nu we de uitleg hebben, kunnen we ook de vergrotende trappen vergelijken.

- Als A *effectiever* (doeltreffender) is dan B, betekent dit dat het doel door A in hogere mate wordt bereikt dan door B.
- Als B *efficiënter* (doelmatiger) is dan A, betekent dit dat het doel door B gemakkelijker wordt bereikt dan door A.

Deze woorden hebben (met een beetje fantasie en goede wil) te maken met twee kernbegrippen in *Metten & Weten*, namelijk *nauwkeurig* en *precies*.

- Een precieze meting die niet nauwkeurig is, is namelijk best wel efficiënt: je voert die normaal gesproken maar één keer uit als je een bepaalde grootte wilt meten.
- Een nauwkeurige meting die niet erg precies is, is daarentegen effectief: je voert die meting namelijk net zo vaak uit totdat je een gemiddelde hebt dat dicht genoeg bij de werkelijke waarde ligt.

Oefening

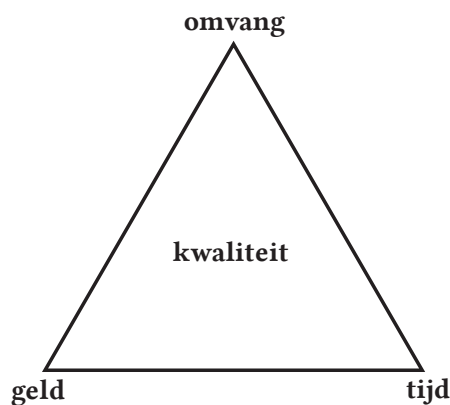
Een hogeschool wil het uitvalpercentage in jaar 1 van een studie verlagen. Het blijkt zo te zijn dat door het stellen van de toelatingseis ‘minstens een 8’ voor Natuurkunde (aanpak A) de uitval afneemt van 55% naar 42%. Een andere aanpak, namelijk het voeren van individuele gesprekken en het afnemen van een toets (aanpak B) de uitval afneemt van 55% naar 41%.

1. Welke van beide aanpakken is *efficiënter*?
2. Welke van beide aanpakken is *effectiever*?
3. Welk gevolg wordt bij het vergelijken van deze twee aanpakken (ten onrechte!) buiten beschouwing gelaten?
4. Kan een aanpak C zowel effectiever als efficiënter zijn?

63. Kwaliteit

Als je met *Google Afbeeldingen* zoekt naar ‘duivelsdriehoek’, kom je niet zomaar in satanische duistere hoeken van het internet terecht, maar wel in plaatjes van driehoeken waar ‘tijd’, ‘geld’ en ‘kwaliteit’ bij staat. De woorden staan niet altijd bij dezelfde hoeken. Sterker nog: ze staan niet altijd bij de hoeken, maar soms bij de zijden. Op *Wikipedia* vind je een plaatje.

EEN PLAATJE DAT NIET EENS EEN VECTORFORMAAT IS. FOEI!
Zou het daarom een duivelsdriehoek heten? Er wordt in de tekst verteld dat het ook wel een *project management triangle* genoemd wordt. Zonder duivel erin dus, maar wel met een extra woord: *scope*. Laten we dat plaatje maar overnemen. En eerst in het Nederlands vertalen. Met het *Linux Libertine* lettertype. En natuurlijk in beleefde vriendelijke kleine letters, WANT HOOFDLETTERS ZIEN ER ZO SCHREEUWERIG UIT!!!



Figuur 63.1 Projectmanagementdriehoek

Het verhaal achter het plaatje is: als je kwaliteit wilt verhogen, zul je meer geld moeten uitgeven. Of extra tijd moeten besteden. Of minder doen (een kleinere *scope*), zodat je meer tijd en geld hebt voor wat je overhoudt.

In de wereld van Agile softwareontwikkeling doen ze iets leuks met dit plaatje. Aan kwaliteit wordt niet getornd: software moet gewoon ‘goed’ zijn volgens een bepaalde norm. Maar wat doe je met die drie hoeken? In de klassieke benadering ligt de scope van een project vast en betekent tegenslag dat je uitloopt (Extra tijd. En ook extra geld). Als je meer mensen inzet, kun je de oorspronkelijke deadline nog steeds halen, maar lopen de kosten nog steeds op. Bij Agile doen ze dat anders. Het aantal beschikbare ontwikkelaars in het team ligt vast en de deadline is hard. Hoe kun je dat bereiken? Door te variëren met de *scope* van het project. Van tevoren weet je niet precies wat je gaat opleveren, maar wel wanneer en tegen welke kosten. Als het tegenzit, krijg je dus minder

functionaliteit. Maar je weet wel wanneer je klaar bent en hoeveel geld ermee gemoeid is. Soms is dat erg prettig.

Die begrippen in die driehoekige plaatjes zijn niet heel lastig. In plaats van ‘geld’ (*money*) zou je ook ‘kosten’ (*cost*) kunnen zeggen. En tijd is tijd (*time*). Het Engelse woord *scope* is iets lastiger. *Van Dale* zegt: gebied, omvang, terrein, sfeer, gezichtsveld, reikwijdte, draagwijdte. Maar we snappen heus wel wat er bedoeld wordt.

‘Kwaliteit’ (*quality*) is het lastigste in dit plaatje. Dat komt doordat het een veel te vaag of te veelomvattend begrip is.

Mocht de lezer een student zijn die nog wil afstuderen (die kans is vrij groot, gezien de doelgroep), let dan tijdens de rest van dit hoofdstuk goed op. Wie weet, krijg je de auteur van dit boek wel als examinerator. Die gaat vervelende vragen stellen als je zegt dat iets ‘beter’ is dan iets anders, of dat je ‘de kwaliteit wilt verhogen’. Het woord ‘kwaliteit’ zegt namelijk weinig. Iets is alleen maar beter dan iets anders als je weet wat je doel is. *Quality is fitness for use*, zeggen de Engelsen. (Sommige dan.)

Een handig hulpmiddel om zinige uitspraken te doen over kwaliteit, is de standaard ISO 25010. Die beschrijft namelijk 31 *kwaliteitskenmerken van software en systemen*. Wanneer is software ‘goed’? Wanneer is het ‘beter’? Het is zinnvoller om te spreken van ‘sneller’, ‘beter onderhoudbaar’, ‘beter leesbaar’, ‘veiliger’, ‘foutbestendiger’.

Metten & Weten gaat over *meten*. (En over *weten*.) Kun je zeggen dat de ene meting *beter* is dan de andere? Ja, maar dan moet je wel weten wat je doel is. De ene meting is sneller, makkelijker, prettiger of goedkoper uit te voeren dan de andere. Nauwkeuriger of preciezer. Minder foutgevoelig en...

Hadden we het in hoofdstuk 0 niet al over WHW? ‘Wat, Hoe en Waarom?’ Inderdaad. Je zult vaak moeten weten *waarom* je iets meet om te weten hoe je het maar beter kunt meten.

Oefeningen

1. Noem tien verschillende kwaliteitsattributen uit de ISO-norm.
2. Hoeveel kwaliteitsattributen bevat die ISO-norm?
3. Welk nummer heeft die ISO-norm?

64. Procenten & Procentpunten

Een blik alcoholarm bier bevat (maximaal) 0,5% alcohol. Hoe hoog zou het alcoholpercentage zijn als het 5% hoger was?

Je bent misschien geneigd om deze vraag te beantwoorden met 5,5%. En het zou inderdaad ook zo kunnen zijn dat de vraagsteller dat bedoelde. Maar eigenlijk verwar je dan de begrippen *procent* (%) en *procentpunt* (pp).

Op Taaladvies.net leggen ze dit als volgt uit.

Een procent is een honderdste deel van iets, bijvoorbeeld: drie procent van tweehonderd is zes. Een procent is dus relatief. Een procentpunt is een punt op een procentschaal en is daarmee een absolute grootheid.

Bijvoorbeeld: een stijging van vier naar vijf is een stijging van één, of een stijging van 25 procent; een stijging van vier naar vijf procent is een stijging van 25 procent of van één procentpunt.

Een percentage kan dus zowel *absoluut* als *relatief* zijn.

De dagelijkse praktijk is weerbarstiger. Als de rente 4,2% is en iemand zegt dat de rente met een half procent gestegen is, ga er dan maar van uit dat hij of zij 0,5 *procentpunt* bedoelt.

Op Taaladvies.net zeggen ze dat er ook omgevingen zijn waarin ze *punt* zeggen als ze *procentpunt* bedoelen. Bij een beursindex bijvoorbeeld. Je leest op een website dat de AEX gestegen is. Er staat namelijk:

AEX 728,79 +2,80 (0,39%)

Dat betekent een stijging van 2,80 procentpunt oftewel 0,39%. De AEX wordt uitgedrukt in een getal dat ooit, op 4 maart 1983, begonnen is op 100 punten. Maar heeft Taaladvies.net dan wel gelijk in de uitspraak dat dit procentpunten zijn? Als de AEX op 100% begonnen is wel. Als de AEX op 100 punten begonnen is niet. En volgens *Wikipedia* waren het eigenlijk geen *punten*, maar *guldens*. Toen de euro op 1 januari 1999 werd ingevoerd, werd die namelijk gelijkgesteld aan 2,20371 gulden en werd de “nieuwe” index door dat getal gedeeld. Dat gebeurde trouwens pas drie dagen later, op maandag 4 januari 1999.

Oefeningen

1. Een blik alcoholarm bier bevat (maximaal) 0,5% alcohol.
Hoe hoog zou het alcoholpercentage zijn als het 5% hoger was?
2. Een blik alcoholarm bier bevat (maximaal) 0,5% alcohol.
Hoe hoog zou het alcoholpercentage zijn als het 5 pp hoger was?
3. Reken de RGB-waarde RGB(0, 0, 174) om in een procentschaal, in zo veel mogelijk decimalen.
4. Bereken met *Excel* in zo veel mogelijk decimalen de kans dat een getal bij een standaardnormale verdeling maximaal één keer de standaardafwijking afwijkt van het gemiddelde.
5. Bereken het verschil tussen je antwoorden op de vorige twee vragen, in procentpunten, in vijftien decimalen.
6. Bereken het verschil uit de vorige vraag in procenten, met twee significante cijfers.

65. Absolute & relatieve onzekerheden

De woorden ‘absoluut’ en ‘relatief’ spelen in het dagelijkse taalgebruik een nogal verschillende rol. Iemand kan ‘absoluut!’ als antwoord geven op een vraag. Dan betekent het ‘ja’ in een sterke vorm. Maar als je zegt: ‘alles is relatief’, dan is er juist weinig sprake van stelligheid.

Als je houdt van zo oorspronkelijk mogelijk Nederlands, zul je in plaats van ‘absoluut’ liever ‘volkomen’, ‘volstrekt’ of ‘zeker’ zeggen. In plaats van ‘relatief’ kies je dan voor ‘betrekkelijk’ of ‘verhoudingsgewijs’. Toch hebben die woorden niet precies dezelfde betekenis. ‘Absoluut’ komt van het Latijnse *absolvere*, dat ‘losmaken’ betekent. Daar zit het wezenlijke verschil met ‘relatief’, waar je iets nou juist niet losmaakt van andere dingen, maar er iets bij betreft of een verhouding aangeeft.

De meetcorrectie die we eerder zagen, bedroeg 3% bij snelheden vanaf 100 km/u. De correctie staat dus niet los van de snelheid, maar hangt ervan af. Je kunt je afvragen of 3% veel of weinig is. Stel dat men in buurland Duitsland een andere norm heeft (wat overigens niet het geval is), dan zou je kunnen zeggen dat het relatief veel of weinig is. Technisch gezien is het interessanter om het te relateren aan de onzekerheid van het meetsysteem. Als die slechts 2% zou zijn, dan is 3% relatief veel. Hoewel... het scheelt maar een factor anderhalf: 3% is 50% meer dan 2%, om het maar eens ingewikkeld te maken. Ook als je iets in verhouding ziet tot iets anders, blijft het betrekkelijk of je het veel of weinig vindt.

Wat nu als je die 3% zou toepassen op een trajectcontrole? De apparatuur stelt vast dat een auto een afstand van 850 meter aflegt in exact één minuut. Wordt de bestuurder van die auto dan bekeurd? Deze kun je uit je hoofd uitrekenen, want er gaan zestig minuten in een uur. De gemiddelde snelheid is dus 51 km/u. Minder dan 3% te hard, maar is 3% niet relatief veel voor een meting over zo’n grote afstand en tijd? Op 850 meter is 3% gelijk aan 25,5 meter en op één minuut is 3% gelijk aan 1,8 seconde. Gevoelsmatig zeg je dat dat wat veel is.

Nog even over beide vormen van meten: die flitspalen stellen de snelheid op één bepaald moment vast, de trajectcontrole geeft een gemiddelde snelheid aan over een wat langere tijd. In geen van beide gevallen is het onmogelijk dat een automobilist op een bepaald moment te hard reed en toch niet bekeurd werd. Het omgekeerde is wel het geval: wie bekeurd wordt, reed sowieso te hard. Absoluut!

Wiskundig zijn absolute en relatieve fouten (relatief) makkelijk weer te geven. Als we een werkelijke waarde x_w willen meten, vinden we iets anders, namelijk de gemeten waarde x_g met een onzekerheid δx .

$$x_w = x_g \pm \delta x$$

De *absolute* onzekerheid is δx , een positief getal (eventueel nul, maar dan zou je een perfecte meting hebben) met (meestal) een eenheid. De *relatieve* onzekerheid δ_{rel} is een dimensieloos getal, meestal uitgedrukt in een percentage.

$$\delta_{rel} = \frac{\delta x}{|x_g|}$$

De absoluutstrepen zijn nodig, omdat de gemeten grootte negatief zou kunnen zijn. Een *relatieve* onzekerheid is dus ook altijd positief.

Eigenlijk zit er iets raars in deze formule. Daarover meer in het volgende hoofdstuk.

Oefeningen

Iemand meet de lengte van een tafelblad en vindt een lengte en een onzekerheid in centimeters. Daarna rekent hij of zij de lengte en onzekerheid om naar *inches*.

1. Wordt de *relatieve fout* door die omrekening groter of kleiner? Of blijft hij gelijk?
2. Wordt de *absolute fout* door die omrekening groter of kleiner? Of blijft hij gelijk?
3. *Die flitspalen stellen de snelheid op één bepaald moment vast, de trajectcontrole geeft een gemiddelde snelheid aan over een wat langere tijd. In geen van beide gevallen is het onmogelijk dat een automobilist op een bepaald moment te hard reed en toch niet bekeurd werd. Leg uit waarom dit zo is.*

66. Kleine waarde, grote onzekerheid

In het vorige hoofdstuk zagen we absoluutstrepen verschijnen voor de relatieve fout. De grootte die je meet, zou namelijk negatief kunnen zijn. Als je een weegschaal zodanig kalibreert dat hij exact 0 g aangeeft als er een leeg bakje op staat, dan zal een meting van datzelfde lege bakje heus wel zoiets kunnen opleveren als $-0,2$ g.

Een spanningsmeting is ook een aardig voorbeeld. Als je een spanning (potentiaalverschil) van ongeveer 0 V zou verwachten, kan het heel goed zijn dat je een negatief aantal millivolts op je scherm krijgt. Niks mis mee: dat is gewoon nul met een onzekerheid. $V_g = -0,01 \pm 2 \text{ mV}$ Mijn absolute fout bedraagt dan 2 mV, de relatieve onzekerheid is... 200%. TWEEHONDERD PROCENT! Is dat geen waardeloze meting dan? Nee. Het heeft alleen geen zin om over relatieve fouten te spreken als je 'nul' verwacht. Alles is relatief veel vergeleken met niks. Weinig is ook al veel als je het afzet tegen niets.

Hoe groot is eigenlijk de relatieve fout als de onzekerheid 2 mV is, de werkelijke waarde 4 mV en de gemeten waarde 5 mV? Dat is een wat vreemde vraag, want we kennen de werkelijke waarde niet. We meten het volgende.

$$V_g = 5 \pm 2 \text{ mV}$$

De relatieve onzekerheid is dus 40%. Maar onze meting wijkt slechts 1 mV af van de werkelijke waarde van 4 mV. Dat is maar 25%. En eigenlijk zou je willen weten hoeveel procent de onzekerheid bedraagt ten opzichte van de *werkelijke* waarde. Hoe onzeker ben ik ten opzichte van wat ik wil weten.

Het probleem is: de werkelijke waarde kennen we niet. Het beste wat we hebben, is de meting, die normaal gesproken aardig in die richting zal liggen. Dit 'probleem' doet zich trouwens alleen maar voor bij onzekerheden die groot zijn ten opzichte van de gemeten waarde. Als je er 100 mV bij optelt, valt het allemaal wel mee. 2 mV op 105 mV is 1,90 % en 2 mV op 104 mV is nog geen promille meer dan dat: 1,92 %.

Terug naar die snelheidsmeting. Er zijn plekken waar je 15 km/u mag rijden. Het voelt wat onrechtvaardig om iemand te bekeuren als hij of zij daar 16 km/u rijdt. Toch is dat 6,7 % te snel, die ene kilometer per uur.

Je ziet trouwens maar zelden flitspalen op plekken waar je maar 15 km/u mag rijden.

Statistiek

67. 'The Average'

Weinig onderwerpen binnen de statistiek zijn zo eenvoudig als het gemiddelde. Je telt gewoon alle getallen bij elkaar op en deelt de uitkomst door het aantal getallen. Het gemiddelde van 6 en 8 is 7.

De Engelse woorden *average* en *mean* zijn niet hetzelfde, al worden ze nogal eens door elkaar heen gebruikt. Salkind (2016) maakt in zijn boek regelmatig gebruik van *Excel* en wijst erop dat dit spreadsheet-programma van Microsoft de verkeerde term gebruikt.

Het is inderdaad gemeen: ze zeggen *average* terwijl ze *mean* bedoelen. En om het nog bonter te maken: 'bedoelen' is in het Engels *to mean* en 'gemeen' is ook al *mean*... Dat kun je toch niet menen!

Een *average* is een maat voor de centrale tendens in een reeks getallen. Dat kan de *mean* zijn, maar ook de *median* (mediaan) of de *mode* (modus). Dat zijn namelijk ook getallen die maatgevend zijn.

Laten we deze drie begrippen eens bekijken aan de hand van een voorbeeld. Voor een tentamen zijn de volgende cijfers gehaald.

1, 4, 4, 5, 5, 6, 6, 6, 6, 6, 7, 7, 7, 7, 8, 8, 8, 9, 9 en 10

Dat zijn twintig resultaten. Als je ze optelt en deelt, kom je op een *gemiddelde* (*mean*) van 6,45.

De *modus* (*mode*) bepaal je door te kijken welk cijfer het vaakst voorkomt. Dat is de 6, het enige cijfer dat vijf keer gehaald is. In dit voorbeeld is de 6 het *modale* cijfer.

De *mediaan* (*median*) is de middelste waarde uit het rijtje als er sprake is van een oneven aantal, en het gemiddelde van de middelste twee als er sprake is van een even aantal. De cijfers staan al op volgorde en het zijn er twintig. Het tiende en elfde cijfer is een 6 en een 7. De mediaan is dus 6,5.

We hebben inmiddels drie *averages* gevonden:

- Het gemiddelde is een 6,45
- De modus is een 6
- De mediaan is 6,5

Stel dat het tentamen anders was uitgevallen en dat dit de cijfers waren geweest.

1, 3, 3, 3, 3, 3, 5, 5, 5, 7, 8, 8, 8, 9, 9, 9, 10, 10, 10, 10

Is dit tentamen nu beter of slechter gemaakt?

Het hoogste en laagste cijfer is hetzelfde. Niemand haalde een 2. Verder zijn er opvallende verschillen. Maar liefst vier tienen en aardig wat negens, maar ook een flink wat mensen met een 3.

We hebben dus drie *averages* gevonden:

- Het gemiddelde is een 6,45
- De modus is een 3
- De mediaan is 7,5

Niet alleen het hoogste en laagste cijfer zijn hetzelfde, maar ook het gemiddelde is exact gelijk aan die van de vorige toets. De mediaan is een heel punt hoger. Maar die modus is nogal schrikken: vorige keer nog een nette 6, nu maar de helft.

De modus is een wat wispelturige maat. Bij twintig tentamencijfers zou je theoretisch de situatie kunnen hebben dat alle cijfers precies twee keer gehaald worden. Dan is elk cijfer een modaal cijfer. Maar stel nu eens dat de 10 maar één keer wordt gehaald en de 9 drie keer: dan is die 9 de modus. Het had ook de 4 kunnen zijn. Of de 1.

De mediaan is wat minder gevoelig voor individuele cijfers. Als de vier studenten die nu een 10 hadden op een 8 waren uitgekomen, zou de mediaan nog steeds hetzelfde geweest zijn.

Het aardige van dit voorbeeld is dat we nog steeds niet gekeken hebben waar de meeste docenten naar kijken: het aantal voldoende (15 tegenover 11). Of liever nog: het percentage geslaagden (75% tegenover 55%).

68. Gemiddelde(n)

Het gemiddelde laat zich het gemakkelijkst uitleggen in woorden: alle getallen optellen en delen door het aantal. Maar het kan ook in formulevorm.

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$$

Hierin is: $\bar{x}$ het gemiddelde, n het aantal waarden en x_i één enkele waarde.

Eigenlijk moet je spreken van het *rekenkundig gemiddelde*. Er is namelijk nog minstens één andere manier om een gemiddelde uit te rekenen: het *harmonisch gemiddelde*.

Wat is het nut van het harmonisch gemiddelde? Stel, je fietst van het gebouw van de Haagse Hogeschool in Delft naar de hoofdvestiging in Den Haag. Je doet dat via de Lange Kleiweg. Dat is – zoals de naam al zegt – niet de kortste route, maar rekent wel handig: volgens Google Maps is het 10,0 km.

Op de heenreis fiets je 20 km/u. Terug heb je tegenwind en haal je maar 15 km/u. Wat is nu je gemiddelde snelheid over de hele fietstocht?

Het lijkt een makkelijk sommetje: het gemiddelde van 20 en 15 is namelijk 17,5. Maar je gemiddelde snelheid blijkt geen 17,5 km/u te zijn.

Op de heenweg zat alles mee. De meewind, maar ook de ronde getallen. Een afstand van 10 km en een snelheid van 20 km/u: dat betekent precies een half uur fietsen. Terug deed je er langer over, namelijk 10/15 (oftewel 2/3) van een uur: 40 minuten.

Je gemiddelde snelheid is de afgelegde weg gedeeld door de tijd. Dat is dus 20 km gedeeld door 30 + 40 = 70 minuten. Omdat we in uren rekenen, moet je delen door 60 van een uur. De uitkomst is $120/7 = 17\frac{1}{7}$ uur. Voor wie liever niet exact rekt: 17,14 km/u. Minder dus dan wat we eerder uitrekenden.

Hoe kan dat nu? Waar is die 0,36 km/u snelheid gebleven? Het punt is: de *afstand* waarover je die 20 km/u haalde was wel gelijk aan de afstand waarover het langzamer ging. Dat was allebei precies 10 km. Maar *de tijd* die je bezig was met langzamer fietsen was wel langer. En dat hakt erin.

Vaak is het begrijpen van dit soort dingen makkelijker als je gaat overdrijven. Bij getallen helpt het soms als je met 0 gaat werken. Wat zou er gebeurd zijn als je op de terugweg écht niet vooruitkwam en met die stevige zuidoostenwind op de heenreis 40 km/u had gehaald. Hoe laat zou je dan thuis

geweest zijn die dag? De rit naar Den Haag toe was een kwartiertje geweest. Maar met 0 km/u zou je nooit meer in Delft komen.

Ook het harmonisch gemiddelde is in een formulevorm te geven.

$$h = \frac{n}{\sum_{i=1}^n \frac{1}{x_i}}$$

En hoewel het formulevormtechnisch minder netjes is, schrijf je ook wel:

$$\frac{1}{h} = \frac{1}{n} \sum_{i=1}^n \frac{1}{x_i}$$

Waarom is dat minder netjes? Omdat je nu niet links van het isteken datgene hebt staan wat je wilt uitrekenen. Een ‘nette’ formule vermeldt links wat je weten wilt en rechts hoe je daaraan komt.

Waarom dan toch die andere formule? Omdat die wat duidelijker laat zien wat het is. Je bepaalt namelijk het gemiddelde van alle inverses en uiteindelijk doe je 1 gedeeld door wat eruit komt.

De *inverse* van x is 1 gedeeld door x oftewel x^{-1} . Je noemt het ook wel *omgekeerde* of de *reciproque* (vaak op z’n Hollands geschreven als ‘reciproke’).

De omgekeerde van een breuk ontstaat door teller en noemer onderling te verwisselen. Let op: het *omgekeerde* is iets anders het *tegengestelde*. Dat krijg je door iets met -1 te vermenigvuldigen.

Naast je gemiddelde snelheid heen en terug is de vervangingsweerstand van twee weerstanden die parallel staan ook een mooi voorbeeld. In serie is het ‘gewoon’ optellen. Parallel is het ‘gewoon’ het harmonisch gemiddelde.

Het *kwadratisch gemiddelde* vind je door het optellen van getallen in het kwadraat en van de som de wortel nemen.

$$x_{RMS} = \sqrt{\frac{1}{n} \sum_{i=1}^n x_i^2} \quad (68.1)$$

De Engelse aanduiding RMS staat voor *Root Mean Square* en dekt de lating beter dan de Nederlandse tegenhanger: het is inderdaad de wortel van het gemiddelde van de kwadraten.

69. Standaardafwijking

De *standaardafwijking* is het kwadratische gemiddelde van de afwijkingen van het gemiddelde.

Eigenlijk is dit hoofdstuk nu al klaar, want meer is het niet. Je zou nog kunnen zeggen dat de standaardafwijking vaak *standaarddeviatie* wordt genoemd en dat het de wortel is van de variantie. Maar in wezen is dat alles.

Laten we nog eens kijken naar die tentamencijfers waarvan we eerder al het gemiddelde, de mediaan en de modus bepaalden.

1, 4, 4, 5, 5, 6, 6, 6, 6, 6, 7, 7, 7, 7, 8, 8, 8, 9, 9 en 10

Het gemiddelde cijfer was een 6,45. Je zou dus van alle behaalde cijfers de afwijkingen van dat gemiddelde kunnen bepalen door het gemiddelde ervan af te trekken. De student met een 10 week 3,55 punt af van het gemiddelde. De studenten met een 6 weken -0,45 punt daarvan af en die ene met een 1 zelfs -5,45. Omdat we het kwadratische gemiddelde willen bepalen, kwadrateren we al die afwijkingen van het gemiddelde. Dat levert grote getallen op voor de uitersten – logisch, want die wijken ook het meest af – en geeft alleen maar positieve getallen, door dat kwadrateren.

Om het met de hand te kunnen uitrekenen en een niet al te grote tabel te krijgen, pakken we vijf cijfers eruit: 4, 5, 7, 7, 9. De modus en de mediaan zie je in één oogopslag: allebei 7. De modus is 7 omdat dat cijfer het vaakst voorkomt. De mediaan is 7 omdat dat het middelste cijfer is als je ze van klein naar groot sorteert.

Het gemiddelde kan ook uit het hoofd: $4+5+7+7+9=32$. Als je dat deelt door 5, kom je op 6,4.

De standaardafwijking wordt al gauw lastig om uit het hoofd te rekenen. Het gemiddelde aftrekken van het cijfer gaat nog wel. Kwadrateren is al pittig. Dan nog het gemiddelde van al die kwadraten... En een wortel trekken uit je blote bol is de meeste mensen ook niet gegeven.

Cijfer	4	5	7	7	9
Afwijking	-2,4	-1,4	0,6	0,6	2,6
Kwadraat	5,76	1,96	0,36	0,36	6,76

Vooruit dan maar, toch de rekenmachine. Of nog makkelijker: *Excel*. De som van die kwadraten is 15,2. Het gemiddelde is die waarde gedeeld door 5 en is dus 3,04. De wortel daarvan is in vier decimalen 1,7436. Omdat we met hele cijfers werken, lijkt het redelijk om één decimaal te nemen: 1,7 of 1,8.

We weten nu hoe we een standaardafwijking berekenen, maar de vraag wat dat getal *is* zou je ook kunnen beantwoorden door het te hebben over de *betekenis* ervan. Je rekent een getal uit, maar wat doe je ermee?

Omdat de standaardafwijking de wortel van de som van de kwadraten van de afwijkingen is, heeft hij altijd dezelfde dimensie als datgene waaraan je rekent. Stel dat je met snelheden werkt: als je het verschil met het gemiddelde neemt blijven het meters per seconde, als je kwadrateert wordt dat m^2/s^2 en als je som neemt blijft het m^2/s^2 . Uiteindelijk neem je de wortel en ben je terug bij snelheden.

In analyses en nadenken over waar je mee bezig bent, is het voor ogen houden van dimensies vaak belangrijk en nuttig.

De vraag was: wat is de betekenis van de standaardafwijking? Het is een maat voor de spreiding van de waarden. Als je voor die oorspronkelijke twintig tentamencijfers de standaardafwijking berekent, kom je (met één decimaal) op 2,0. Hoe groot zou die voor dat tweede tentamen zijn? Meer of minder dan 2,0? Laten we de resultaten er nog eens bij pakken.

1, 3, 3, 3, 3, 3, 5, 5, 5, 7, 8, 8, 8, 9, 9, 9, 10, 10, 10, 10

In de eerste toets was er ook een 1, maar er waren geen tweeën en drieën. En er was maar één 10. Bovendien waren er meer cijfers rond de 6 en 7. Veel cijfers die ver afwijken van het gemiddelde dus. En weinig die in de buurt van het gemiddelde zitten. Zonder het uit te rekenen kun je ‘aanvoelen’ (eigenlijk: beredeneren) dat de standaardafwijking groter zal zijn. Als je het uitrekent, blijkt dat te kloppen: 2,94. Bijna een punt meer.

70. Een klein aantal metingen

In het dagelijks leven gaan mensen er gemakkelijk van uit dat metingen domweg kloppen en geen onzekerheid in zich hebben. Iemand kijkt op een thermometer en ziet dat het $22\text{ }^{\circ}\text{C}$ is. Dat die thermometer aan een koude muur hangt of in de volle zon, niet geijkt is... het zal allemaal wel.

Hoe zit dat dan als je één meting hebt? Die $22\text{ }^{\circ}\text{C}$, kun je daar statistiek op bedrijven? De modus, de mediaan en het gemiddelde bestaan wel voor één meting, maar ze zijn onzinnig. Of beter gezegd: ze voegen niets toe aan wat je al hebt. Ze zijn voor de ‘metingenreeks’ in dit geval alle drie gelijk aan $22\text{ }^{\circ}\text{C}$.

Hoe zit dat met de standaardafwijking? Is die gedefinieerd voor één waarde? Het antwoord is ja, maar de uitkomst is nog gekker: de afwijking tussen één meting en het gemiddelde is gelijk aan 0. Nee, dat is niet waar: er moet nog een eenheid bij. In ons geval is de standaardafwijking exact gelijk aan $0\text{ }^{\circ}\text{C}$.

Is het raar om één waarde als beste schatting aan te nemen als je maar één meting hebt? Nee. Het is geen al te beste schatting, maar je hebt niks beters. Bovendien: vaak slaat het ook wel ergens op. Die thermometer maakt erg aannemelijk dat het niet vriest. Maar het zou in werkelijkheid natuurlijk best $15\text{ }^{\circ}\text{C}$ kunnen zijn als die thermometer in de volle zon hangt.

Je zou de meting natuurlijk meteen kunnen herhalen en dan lees je waarschijnlijk hetzelfde af. Of een uur later nog een keer en dan zien dat hij op $24\text{ }^{\circ}\text{C}$ staat. Hebben we dan een spreiding van $2\text{ }^{\circ}\text{C}$ in de meting? Nee, we hebben twee verschillende metingen gedaan. We vergelijken dan de temperatuur om (bijvoorbeeld) half een ’s middags met de temperatuur om half twee ’s middags.

Misschien is de temperatuurmeting wat serieuzer van aard. Niet een huistuin-en-keukenthermometer die toevallig in de zon hangt en een uur later opnieuw wordt afgelezen, maar een geijkte sensor in een vat met een bepaalde vloeistof. We meten heel goed en met verschillende sensoren op hetzelfde moment, en toch zit er een spreiding in de waarden die je vindt. Er zit dan een onzekerheid in de temperatuur die een gevolg is van datgene waaraan je meet. Een definitieprobleem dat zorgt dat de onzekerheid van het meetinstrument of de methode niet het eigenlijke euvel is.

Wat als we twee metingen hebben? Het gemiddelde van $22\text{ }^{\circ}\text{C}$ en $24\text{ }^{\circ}\text{C}$ is $23\text{ }^{\circ}\text{C}$. Is het zinnig om dat als beste schatting te nemen? Ja. Is er een standaardafwijking te berekenen? Dat is lastiger te zeggen.

Beide metingen wijken 1 °C af van het gemiddelde. De standaardafwijking vind je door ze kwadratisch op te tellen. Maar wat doe je dan eigenlijk? Je telt twee keer hetzelfde getal bij zichzelf op. Wiskundig kun je dat algemeen opschrijven.

$$\bar{x} = \frac{x_1 + x_2}{2}$$

$$\Delta x_1 = |x_1 - \bar{x}| = \left| x_1 - \frac{x_1 + x_2}{2} \right| = \left| \frac{x_1 - x_2}{2} \right|$$

$$\Delta x_2 = |x_2 - \bar{x}| = \left| x_2 - \frac{x_1 + x_2}{2} \right| = \left| \frac{x_2 - x_1}{2} \right|$$

De verschillen zijn aan elkaar gelijk. Twee metingen wijken altijd evenveel af van hun gemiddelde. Logisch, want dat zit precies midden tussen die twee waarden in.

De standaardafwijking van twee getallen is gelijk aan de kwadratische som van hun afwijkingen, die allebei gelijk zijn aan dezelfde waarde Δx .

$$\sigma = \sqrt{\frac{2(\Delta x)^2}{2}} = \Delta x$$

Op die manier bereken je dus een spreiding door het (kwadratische) gemiddelde te nemen van de afwijkingen. Maar het is helemaal geen gemiddelde. Het is gewoon de afwijking zelf. Er wordt niet opgeteld en gedeeld, want er is er maar eentje.

Is dat een vreemde keuze als waarde? Nee. Het is niet vreemd om op grond van die twee metingen te zeggen dat de temperatuur gelijk is aan:

$$T = (23 \pm 1) \text{ °C}$$

Moet daar niet een significant cijfer bij? We zouden achter een 1 toch altijd nog een cijfer hebben? Nee, niet altijd. Dat doen we alleen maar als er een cijfer voorhanden is dat significant is. Bijvoorbeeld als je het gemiddelde zou berekenen van de volgende veertien metingen: 22 °C, 24 °C, 23 °C, 23 °C, 23 °C, 24 °C, 24 °C, 23 °C, 21 °C, 23 °C, 24 °C, 23 °C, 24 °C, 21 °C. De som van deze metingen is 322 °C en als je dat deelt door 14, kom je op exact 23 °C. Alle afwijkingen zijn gelijk aan 0, 1 of 2 °C en de som van de kwadraten daarvan is ‘toevallig’ precies 14 °C². (Let op de eenheid: dit zijn graden Celsius *in het kwadraat*.) Daardoor komt de standaardafwijking op exact 1 °C uit.

Als de laatste meting 22 °C was geweest, dan zou het resultaat anders luiden: de som was dan een klein beetje meer, de spreiding een klein beetje minder.

$$T = (23,1 \pm 0,9) \text{ °C}$$

Daarom was het in die reeks van veertien metingen met 21 °C als laatste wél legitiem om een extra cijfer te geven.

$$T = (23,0 \pm 1,0) \text{ °C}$$

Terug naar die standaardafwijking van twee metingen. Er is op zich niet heel veel mis mee om die toch maar gewoon zo uit te rekenen of domweg het (halve) verschil van twee waarden te nemen. Voor het begrip is het goed om te zien dat de informatie die je krijgt uit weinig metingen beperkt is.

Met één meting kun je wel een beste schatting doen van het gemiddelde (namelijk die ene waarde), maar je kunt *helemaal niets* zeggen over de spreiding.

Met twee metingen kun je een iets betere beste schatting maken van het gemiddelde (namelijk het gemiddelde van die twee), maar je kunt *maar heel weinig* zeggen over de spreiding (namelijk het verschil tussen die twee).

Wat gebeurt er met het gemiddelde als er een meting bij komt? En wat met de standaardafwijking? Dat weet je niet. Dat hangt af van de toevallige waarde van die ene meting. In ons (gekunstelde) voorbeeld waar het gemiddelde van veertien metingen exact 23 °C bleek te zijn en de standaardafwijking 1 °C, zou je drie metingen van 23 °C kunnen toevoegen. Het gemiddelde blijft dan gelijk en de standaardafwijking wordt (afgerond) 0,9 °C. Zou je in plaats van drie keer 23 °C twee keer 20 °C en één keer 29 °C toegevoegd hebben, dan bleef het gemiddelde ook exact 23 °C en werd de standaardafwijking verdubbeld: er zou dan exact 2 °C uitgekomen zijn. In correcte notatie:

$$T = (23 \pm 2) \text{ °C}$$

Vreemd, maar volgens de regels. Het is net zo exact als die eerdere berekening, maar we geven een decimaal minder op. Het aardige van deze (gekozen) waarden is dat de ene 23 °C de andere niet is. Het gemiddelde van de eerste twee was toevallig 23 °C. De daaropvolgende drie waren allemaal gelijk aan diezelfde waarde. Hadden we niet de eerste en de tweede, maar de laatste en de tweede waarde gemeten, dan was het gemiddelde gelijk geweest aan 26,5 °C. Het verschil tussen 24 °C en 29 °C is 5 °C, dus de spreiding was dan 2,5 °C geweest.

Oefeningen

1. Maak een werkblad in *Excel* en neem deze getallen over in kolom A.

22 24 23 23 23 24 24 23 21 23 24 23 24 21

2. Neem in kolom B in je werkblad formules op waarmee je het gemiddelde van de getallen berekent. Stel de formule zodanig op dat je het aantal elementen in kolom A kunt veranderen, door gebruik te maken van een formule waarin de hele kolom wordt meegenomen.

=GEMIDDELDE(A:A)

3. Voeg in kolom B in je werkblad formules toe waarmee je onderstaande uitkomsten kunt berekenen.
 - a. de som
 - b. het aantal waarden
 - c. het gemiddelde (som gedeeld door aantal waarden)
 - d. de mediaan
 - e. de modus
 - f. de grootste waarde
 - g. de kleinste waarde
 - h. de grootste afwijking van het gemiddelde
 - i. de gemiddelde absolute afwijking
 - j. het aantal waarden dat groter is dan het gemiddelde
4. Bedenk wat er verandert aan de tien uitkomsten in kolom B als je het getal 23 zou toevoegen aan de veertien waarden die je al hebt.
5. Bedenk wat er verandert aan de tien uitkomsten in kolom B als je het getal 20 zou toevoegen aan de veertien waarden die je al hebt.
6. Test of wat je bedacht hebt in de vorige twee oefeningen klopt door die waarden toe te voegen aan kolom A. (Let op: in beide gevallen doe je een berekening op vijftien waarden.)
7. Als de kans dat iemand iets heeft of doet één op miljard is, hoe groot is dan de kans dat iemand op de wereld dat heeft of doet?

71. Eindig aantal mogelijkheden

Aan het tentamen hadden twintig studenten deelgenomen en je kon alleen maar hele cijfers halen. Het aantal mogelijke uitkomsten was dus eindig. Er waren veel mogelijke cijferlijstjes denkbaar, maar dat hield wel een keer op.

Elke student had tien mogelijke cijfers om te halen, want er stonden geen cijfers achter de komma. Hoeveel mogelijkheden zijn dat? Voor die eerste student tien, voor de tweede ook. Voor hen samen dus honderd mogelijkheden: naast elk cijfer dat student A haalde, waren er voor B tien mogelijkheden.

Als het tentamen nu eens met ‘voldoende’ of ‘onvoldoende’ was beoordeeld? Dan waren er twee mogelijke ‘cijfers’ geweest. Bij twee studenten zijn er dan vier mogelijkheden. Bij twintig studenten worden dat er $2^{20} = 1.048.576$. Ruim een miljoen. Best veel, maar wel een *eindig* aantal. Bij 33 studenten zou je al meer mogelijkheden hebben dan er mensen op de wereld zijn.

Als de toets door 100 studenten gedaan was, zou het aantal mogelijke cijferlijsten 2^{100} zijn geweest. Dat is 1.267.650.600.228.229.401.496.703.205.376, wat je liever schrijft als een afgeronde macht van 10: $1,268 \cdot 10^{30}$. Een getal met dertig cijfers... Nee, met 31 cijfers, want 10 tot de macht N heeft N+1 cijfers. Ga maar na: $10^2 = 100$ en dat zijn drie cijfers.

Bij 150 studenten krijg je een getal dat niet meer op één regel past en als je alle cijfers van alle studenten die ooit een bepaalde studie deden zou verzamelen, praat je al gauw over een paar duizend: getallen met pakweg 300 of 400 cijfers. Ja, dat zijn flinke jongens. Maar het aantal blijft eindig.

We hadden het over de situatie waar je voldoende of onvoldoende scoort. Met tien mogelijke cijfers gaat het veel sneller. Wat is eigenlijk het grootste getal dat nog op een regel past?

10.000.000.000.000.000.000.000.000.000.000.000.000.000.000.000.000.000.000

Dat is 10 tot de macht 58. Hoeveel studenten heb je nodig om zoveel mogelijkheden te hebben? Inderdaad: 58. En hoeveel studenten zorgen samen voor dit aantal als je toch weer die o/v-score invoert? Dan moet je de logaritme nemen van het getal. De ${}^2\log$ om precies te zijn. Hoe moest dat ook alweer?

De $\log_g$ van x reken je uit door $\log_{10}(x)$ te delen door $\log_{10}(g)$. Bereken dus twee keer de $\log_{10}$ op je rekenmachine en deel de uitkomsten. ${}^{10}\log(10^{58}) = 58$ en $\log_{10}(2) \approx 0,3010$. Delen en afronden op een geheel getal naar boven(!) leert ons dat je met 1871 studenten op dit aantal mogelijkheden zit.

Stel nu eens dat we cijfers zouden geven die niet in gehele getallen werden uitgedrukt, maar één decimaal hadden. Dat je ook een 5,4 kon halen in plaats van een 5. Of een 9,9 in plaats van een 10.

Als je hele cijfers geeft, heb je tien mogelijkheden. Heb je er met één decimaal dan honderd? Nee. Als je een 10 haalt, is het cijfer achter de komma altijd een 0, omdat er hoger niet gegeven wordt. De cijfers 1 t/m 9 kunnen tien verschillende ‘varianten’ hebben, maar een 10 is een altijd een 10. Er zijn dus in totaal $9 \times 10 + 1 = 91$ mogelijkheden.

Wie heeft ooit verzonnen om van 1 t/m 10 te tellen? Met een tientallig stelsel is 0 t/m 9 is veel logischer. Bovendien is 10 geen *cijfer*, maar een *getal*. Nog onlogischer wordt het als je bedenkt dat 1 het minimum is. Wie niets invult op het antwoordvel of alle meerkeuzevragen fout doet, zou ook niets (niks, nul, noppes) als waardering moeten krijgen.

Als je van 0 t/m 9 zou tellen, zou een cijfer achter de komma wél tot honderd opties leiden. Hierbij is er één probleem: je moet dan wel accepteren dat je alle cijfers naar beneden afrondt. Iemand met een 9,9 (alles goed!) zou in hele cijfers dan een 9 hebben. Die voelt zich ook bekocht.

Hoeveel mogelijke uitslaglijstjes heb je met 120 studenten en 91 mogelijke cijfers? Inderdaad: 91^{120} . Dat is $1,2161 \cdot 10^{235}$. Ja, dat is een groot getal. Hoeveel ruimte heb je nodig om al die mogelijke lijstjes op te slaan. Laten we zeggen dat je één byte nodig hebt voor een cijfer, Vooruit: een halve byte per cijfer, want met vier bits kun je ook wel tien (hele) cijfers opslaan. (Zestien zelfs.) Uitgaande van een harde schijf van één miljard terabyte heb je dan nog steeds $6,1 \cdot 10^{126}$ van die schijven nodig. Sterker nog: als elke wereldbewoner zo’n schijf ter beschikking zou stellen, red je het nog niet. De exponent daalt met negen of tien, maar dat zet geen zoden aan de dijk. Al zou elke ster in het heelal (naar schatting zo’n 10^{22}) een miljard planeten bevatten waarop alle bewoners hun kolossale harde schijven ter beschikking stelden: het schiet nog steeds tekort.

Sommige getallen zijn domweg te groot om te bevatten. Hoeveel energie hebben we eigenlijk nodig voor zoveel computers? De zon levert een vermogen van zo’n $4 \cdot 10^{22}$ Watt. Dat is aardig wat.

Is het aantal mogelijke uitkomsten van een toets met 120 studenten en 91 mogelijke cijfers oneindig? Nee. Erg veel, laten we het daar op houden.

Oefening

Hoe zou je op een beknopte manier tentamencijfers in een database opslaan?

72. Kansvariabele

In de kansrekening is een stochastische variabele of stochastische grootheid een grootheid waarvan de waarde een reëel getal is dat afhangt van de toevallige uitkomst in een kansexperiment.

De stochastische variabele, ook toevalsvariabele, kansvariabele of stochast, is een eigenschap van de uitkomst die in een getal is uit te drukken, zoals de leeftijd of het inkomen van een toevallige voorbijganger. Het toeval bepaalt de uitkomst van het experiment, en bijgevolg is de waargenomen waarde van de stochastische variabele ook afhankelijk van het toeval. Bij een onderzoek naar de verdeling van de leeftijd is niet de toevallige voorbijganger zelf, de uitkomst, van belang, maar z'n leeftijd, een eigenschap van de uitkomst. Die leeftijd is in dit geval de stochastische variabele. Zo zijn ook het inkomen van een willekeurig gekozen Nederlander en het aantal keren dat 'kruis' gegooit wordt in een serie van 100 worpen met een munt, stochastische variabelen. Hoewel voor elke mogelijke uitkomst de waarde van de stochastische variabele vastligt, hangt de waargenomen waarde af van het toeval, als gevolg van de toevallige uitkomst. Formeel is een stochastische variabele daarmee een functie die aan elke uitkomst een getal, de waarde van de bedoelde eigenschap, toevoegt.

Bron: Wikipedia

De uitleg op *Wikipedia* van het begrip kansvariabele beslaat een halve pagina in dit boek. Is het echt zo'n lastig begrip? Of is het alleen maar moeilijk voor mensen die het niet begrijpen?

Het lijkt een beetje op de begrippen *object* en *klasse* in objectgeoriënteerd programmeren. Is 'fiets' een object of een klasse? Dat ligt er maar aan. Als je 'de fiets in z'n algemeenheid' bedoelt, dan is het een *klasse*. Als je die ene fiets bij jou thuis in de schuur bedoelt, is het een *object*.

"De fiets staat in de belangstelling van veel mensen."

Is 'de fiets' in deze zin een object of een klasse? Waarschijnlijk een klasse, want we bedoelen het *verschijnsel*, het *begrip*. Het *fenomeen*, het 'algemene ding'. Niet één bepaalde fiets. Tenzij je natuurlijk wél één heel bepaalde fiets bedoelt.

Een gek voorbeeld. In een woonwijk wordt er om onverklaarbare redenen telkens een fiets bij iemand in de tuin aangetroffen. Telkens weer bij een ander huis, niemand weet van wie die fiets is of wie dat ding daar neerzet. Het gaat al weken zo en op een gegeven moment is De Fiets een onderwerp van gesprek. Die ene fiets dus, dat rare stuk schroot op twee wielen dat overal

opduikt. In dat geval bedoel je met ‘De Fiets’ die in de belangstelling staat toch een specifiek voorbeeld, een heel concrete instantie van de klasse ‘fiets’.

Hou dit in je achterhoofd: er is een verschil tussen het algemene (abstracte) geval en dat ene specifieke (concrete) geval.

Nu terug naar de kansvariabele en een voorbeeld van *Wikipedia*.

Bij het gooien met twee dobbelstenen bestaat de uitkomstenruimte uit de $6^2 = 36$ paren mogelijke ogenaantallen:

$$\begin{aligned}\Omega &= \{1, 2, 3, 4, 5, 6\} \times \{1, 2, 3, 4, 5, 6\} \\ &= \{(1, 1), (1, 2), \dots, (1, 6), (2, 1), \dots, (6, 6)\}\end{aligned}$$

Het totale aantal geworpen ogen wordt gedefinieerd door de stochastische variabele:

$$X(\omega_1, \omega_2) = \omega_1 + \omega_2$$

Het waardenbereik van X is $\{2, 3, \dots, 12\}$.

Het totale aantal geworpen ogen is dus de stochastische variabele. Dat zijn elf (geen twaalf!) verschillende uitkomsten. Elf verschillende reële getallen, die allemaal een eigen kans van voorkomen hebben.

Toch kun je ook zeggen dat *het totale aantal geworpen ogen* géén stochastische variabele is, als je er iets anders mee bedoelt. In een klas zitten twaalf studenten. We willen door loting een bepaalde student (“Student 1”) koppelen aan een willekeurige andere student (“Student 2” t/m “Student 12”). Dat doen we door met twee dobbelstenen te gooien. Nadat we gegooid hebben, liggen de twee stenen op tafel. We zijn benieuwd wie het geworden is!

Het totale aantal geworpen ogen blijkt 7 te zijn.

Is *het totale aantal geworpen ogen* nu nog steeds een kansvariabele? Nee. Als er eenmaal gegooid is, zijn er geen kansen meer. Het is zoals het is. Je zou hooguit nog kunnen zeggen dat het kansvariabele is waarbij de uitkomst 7 een kans van 1 heeft en alle andere uitkomsten een kans van 0.

Iets ingewikkelder wordt het als die twee dobbelstenen in een niet-doorzichtige beker zitten. We schudden de beker en zetten die met een behendige zwaai – zonder dat de dobbelstenen eruit vallen – omgekeerd op tafel.

Is *het totale aantal geworpen ogen* nu een kansvariabele? Of toch niet?

Je ziet aan dit voorbeeld dat het begrip kans op zich wat lastig en abstract is. Dat *totale aantal geworpen ogen* ligt na het omkeren van die beker gewoon vast! Je kunt wel hopen of bidden dat het 7 is, maar die stenen liggen daar. Het lot is geworpen en we richten alle ogen op de dobbelsteen. We *weten*

alleen het aantal nog niet. De *kans* dat je goed zo *raden* hoeveel het er zijn, zou je weer wél als een stochastische variabele kunnen zien.

“Je kunt wel hopen of bidden dat het 7 is...”

Dit is nog een aspect van begrippen als ‘toeval’ en ‘kans’. Bestaat er wel zoiets als toeval? Is het echt willekeur wat er uit die worp komt, of beschikken God, je Karma, je Sterrenbeeld of Het Universum daarover?

Veel mensen denken dat die dobbelstenen toeval zijn, net als kop of munt. Bij het krijgen van een ernstig ongeluk of het tegenkomen van een bijzonder persoon op een bijzonder moment ervaren sommige van die mensen dat anders. “Dat kan geen toeval zijn!”

Als 60% van de studenten elk jaar het tentamen in één keer haalt, heb jij dan een slagingskans van 0,6? Of ben jij die ene die altijd alles (of nooit iets) in één keer haalt?

Ook als je niet in de invloed van God en Het Universum gelooft is dat kansbegrip lastig. Die dobbelstenen vallen niet *toevallig* zo. Dat hangt af van hoe je gooit en hoe de lucht stroomt en tot op microscopisch niveau de massa binnen de dobbelstenen verdeeld is. Toch hoeft je niet te discussiëren over levensbeschouwelijke onderwerpen en je visie of dingen toeval zijn of niet om met kansen te rekenen. Dat je die ene bijzondere persoon op een bijzonder moment tegenkwam, is een eenmalige gebeurtenis. Wat dat toeval? Was het voorbestemming? Statistiek leent zich niet voor eenmalige gebeurtenissen. Maar als je duizend keer met een dobbelsteen gooit, zul je zien dat het gemiddelde echt niet ver van de 3,5 ligt. En als je het honderdduizend keer doet, wordt dat nog duidelijker.

73. Verwachtingswaarde

Iemand vraagt of je lootjes wilt kopen. De opbrengst is voor een goed doel. Er worden er honderd verkocht en ze kosten één euro per stuk. Valt er iets te winnen dan? Ja! Eén prijs van € 50, drie prijzen van € 20 en vijf van € 10.

Vraag 1: Is het een goed idee om lootjes te kopen?

Vraag 2: Hoeveel dan?

Als het antwoord op vraag 1 ‘nee’ was, is het antwoord op vraag 2: ‘nul’. Eigenlijk is het in dat geval maar één vraag dus.

Om tot je keuze te komen, kun je gebruikmaken van de *verwachtingswaarde*.

De verwachtingswaarde van een stochastische variabele is de waarde die deze variabele ‘gemiddeld genomen’ zal aannemen. Dit gemiddelde is het gewogen gemiddelde van alle mogelijke uitkomsten met als gewichtsfactor de kans dat een bepaalde waarde zich voordoet.

Pascal & Fermat bedachten in 1654 dit begrip bij hun oplossing van het zogenaamde *puntenprobleem*. Twee partijen spelen een spel waarbij je een gelijke kans hebt om te winnen. De partij die als eerste zes keer heeft gewonnen, wint de pot met zestig muntstukken. Stel dat het spel afgebroken moet worden bij een stand van 5-3, hoe moet de pot dan verdeeld worden? Oplossing: zie *Wikipedia* (‘puntenprobleem’).

Bij die loterij is er een kans van 1% dat je die € 50 wint, 3% dat je € 20 wint en 5% dat je € 10 wint. De verwachtingswaarde van de ‘winst’ bij één lot dus:

$$0,01 \times € 50 + 0,03 \times € 20 + 0,05 \times € 10 = € 0,50 + € 0,60 + € 0,50 = € 1,60$$

Dat is niet de *winst*, maar de *opbrengst*. Om de winst te berekenen, moet je de *kosten* (€ 1) ervan aftrekken. De verwachtingswaarde van de winst is € 0,60.

Is de kans dat je geld wint groter dan de kans dat je geld verliest? Nee! De kans dat je iets wint is namelijk $1\% + 3\% + 5\% = 9\%$ en de kans dat je verliest is dus 91%. Maar als je verliest, verliest je niet veel. Als je wint, win je méér.

Er is overigens wel nog een probleem. De opbrengst is voor een goed doel... Als jouw kans op winst groter is, is de kans op verlies voor dat goede doel kleiner. Je kunt dus beter geen loten kopen en een tientje overmaken aan dat goede doel. Een definitieprobleem? Nee. De begrippen zijn wel duidelijk. We hebben alleen het *doel* van het kopen van de lootjes niet bepaald.

Als je het doel niet goed kent, weet je niet wat je het beste kunt doen.

Iemand laat twee enveloppen zien waar geld in zit. Je mag er eentje hebben. Het enige wat je weet, is dat in de ene envelop *twee keer zoveel* zit als in de andere. Je maakt de envelop open: € 10. ‘Wil je nog ruilen?’

Je denkt even na. In die andere envelop zit óf de helft (€ 5) óf het dubbele (€ 20). Beide mogelijkheden zijn even waarschijnlijk (50%).

De verwachtingswaarde van de inhoud van de envelop die je krijgt door te ruilen is dus € 12,50... Toch?

Maar wat nou als je die andere als eerste had genomen?

De enveloppenparadox.

Oefeningen

1. Wat is de verwachtingswaarde van het kwadraat van het totaal aantal ogen dat je gooit met twee dobbelstenen?

2. Waar zit de redeneerfout in de enveloppenparadox?

De oefeningen hieronder gaan uit van de in dit hoofdstuk genoemde loterij.

3. Bekijk twee verschillende situaties.
 - A. Je doet twee keer mee met de loterij.
 - B. Je doet één keer mee en je koopt twee lootjes.Wanneer is de verwachtingswaarde van de winst groter?
4. Wat is de verwachtingswaarde van de winst bij vijftig loten?
5. Bij welk aantal loten is de winst geen stochastische variabele?
6. Bekijk twee verschillende situaties.
 - A. Je doet twee keer mee met de loterij en koopt beide keren tachtig lootjes.
 - B. Je doet één keer mee en je koopt 92 lootjes.Wanneer is de verwachtingswaarde van de winst groter?
Wanneer is de kans dat je een prijs wint groter?
7. Bedenk argumenten om wel of niet mee te doen met de *Nationale Postcode Loterij*.

74. Permutaties & Combinaties

De filosoof en wetenschapper Leibniz schreef in 1714 in een brief dat het even makkelijk is 11 als 12 te gooien met twee dobbelstenen. Er was namelijk volgens hem maar één manier om elk van beide getallen te gooien.

Het klopt dat je 11 met twee dobbelstenen alleen maar kunt gooien met een 5 en 6. En het klopt ook dat 12 alleen maar kan met twee zessen. Maar wat Leibniz – toch echt geen domme jongen – over het hoofd zag, is dat er *twee* manieren zijn om 5 en 6 te gooien: eerst een 5 en dan een 6, of eerst een 6 en dan een 5.

Wanneer zou je die denkfout niet maken? Misschien als je heel goed nadacht over alle mogelijkheden. Met één dobbelsteen heb je zes mogelijke uitkomsten. Met twee dobbelstenen heb je voor elke van die zes mogelijke uitkomsten nog eens een keer zes mogelijke uitkomsten. Je hebt er dus $6 \times 6 = 36$. En als je die alle 36 in een tabelletje zet, zie je ook dat niet alleen 5 en 6, maar ook 6 en 5 de gevraagde 11 oplevert.

Je ziet dan ook meteen dat 7 de meest waarschijnlijke uitkomst is. Die kun je namelijk vormen uit 1+6, 2+5, 3+4, 4+3, 5+2 en 6+1. De kans om 7 te gooien is maar liefst zes keer zo groot als de kans op 12.

Dat Leibniz zich vergiste, is menselijk. Misschien heeft hij heel hard gelachen toen iemand hem op die domme fout wees. Het zou ook een grap geweest kunnen zijn trouwens. Er zijn mensen die beweren dat de kans op het winnen van de Staatsloterij 50% is, om het simpele feit dat er maar twee mogelijke uitkomsten zijn. Je wint 'm wel of je wint 'm niet.

Sommige mensen menen dat laatste nog ook... Er zijn twee mogelijkheden, dus 50% kans. Omdat je zo lang mogelijk de optie open moet houden dat iemand anders gelijk heeft: het zou natuurlijk kunnen zijn dat die mensen een andere definitie van het begrip 'kans' hebben.

Laten we het erop houden dat de vergissing van Leibniz bewijst dat statistiek moeilijk is. Altijd handig als je zelf de fout in gaat. 'Ken je dat verhaal van Leibniz?'

Dat iets moeilijk is, maakt het makkelijker om het jezelf niet al te kwalijk te nemen dat je het fout doet. Daar schiet je sowieso niet veel mee op. Maar met *voorkomen* dat je fouten maakt en/of ze zo snel mogelijk vinden, schiet je wél iets op.

Bij statistische dingen kan het helpen dat je een simulatie schrijft met een computerprogramma. Bijvoorbeeld in *Python*. Even zoeken met Google en je

weet hoe je een *random number* genereert. In dit geval: een geheel willekeurig getal van 1 t/m 6, want het is een dobbelsteen.

Zou dit 'm zijn?

```
import random

dobbelsteen1 = random.randint(1,6)

print(dobbelsteen1)
```

Nee! Als je dit programma test, lijkt het werken. Er komt 3 uit of zo. Of 1. Of 5, of 4, of nog een keer 3. Maar nooit 6, want het tweede argument van `randint()` is de bovengrens en die wordt nooit bereikt. En om dat aan te tonen, doe ik het honderd keer.

```
2 5 5 1 1 3 1 1 3 4 5 5 2 6 3 6 2 1 2 1 4 2 2 3 5 4 5 4 6 3 2 6 2 2 1 1 6 4 5 3 5 6 5 4
1 4 3 2 4 4 4 6 1 6 1 6 3 4 2 2 4 5 6 1 6 1 1 5 4 1 3 5 6 4 3 6 3 3 2 1 4 2 3 2 3 1 4 3
4 3 2 1 1 5 1 4 3 4 1 4
```

Hé, daar staat dus toch 6 tussen... Waarom dacht ik dan van niet? Omdat in sommige programmeertalen en functies die bovengrens niet meedoet. Neem het Java-programma hieronder. Als je denkt dat die 6 ervoor zorgt dat de opties 1 t/m 6 voorkomen, dan heb je het mis. En dan nog iets: deze code kan ook de waarde 0 opleveren. Er zijn wel zes mogelijke uitkomsten, maar dat zijn dus de getallen 0 t/m 5. Er is niet veel voor nodig om het correct te krijgen. Maar dat geldt voor zo veel bugs.

```
public static void main( String args[] ) {
    Random rand = new Random();
    int upperbound = 6;
    using nextInt()
    int int_random = rand.nextInt(upperbound);
    System.out.println("Random integer value from 0 to" +
        (upperbound-1) + " : " + int_random);
}
```

Hoe voorkom je nu deze fout? Geloof het of niet: deze vijf dingen helpen.

1. De specificaties van de functie lezen.
2. Goed nadenken van tevoren.
3. Je concentreren tijdens het schrijven van de code.
4. Heel veel testen op een slimme manier.
5. Iemand anders ernaar laten kijken.

Dat ik het programma honderd keer uitvoerde, was geen bewijs dat het goed werkte. Maar ik had aan de uitvoer wel een aantal dingen gezien die me een beetje geruststelden.

1. Het getal 0 kwam nooit voor
2. Het getal 1 wel
3. Het getal 6 ook
4. Alle getallen daartussenin ook
5. Geen enkel getal opvallend vaak

Dat vierde geloofde ik sowieso wel, maar toch goed om te zien.

Nu ik weet hoe je willekeurige getallen maakt in *Python*, kan ik mijn programma schrijven. Ik kom op deze code.

```
import random
import numpy as np

uitkomst = np.zeros(13)

for i in range(100):
    dobbelsteen1 = random.randint(1,6)
    dobbelsteen2 = random.randint(1,6)
    worp = dobbelsteen1 + dobbelsteen2
    uitkomst[worp] = uitkomst[worp] + 1

print(uitkomst)
```

Is het algoritme goed? Nadenken helpt, hadden we gezien. Dat ik met een array van *dertien* nullen begin (en niet met twaalf) is niet om te bewijzen dat ik niet bijgelovig ben, maar omdat de hoogste index 12 moet zijn. Bovendien wil ik zien dat 0 nooit voorkomt.

Ik test het programma en kom op deze uitvoer.

```
[ 0.  0.  0.  2.  9. 14.  9. 17. 18. 13. 14.  4.  0.]
```

Tsja. Wat moet je ervan zeggen? Het ‘valt op’ dat de waarden aan de rand niet of weinig voorkomen en in het midden veel meer. Dat had ik al verwacht in het verhaal met die vele manieren om 7 te gooien. Maar bewijst dit nu iets?

Nee. Iets bewijzen doet een simulatie zelden of nooit. Maar waarom herhaal ik het eigenlijk 100 keer en geen 1000 keer, of 10000 keer of 100.000 keer? Zelfs een miljoen keer kan mijn computer nog wel in een paar seconden aan. Tien miljoen lukt ook nog binnen een minuut. En ik heb toch trek in koffie...

De uitkomst is er! En de koffie wordt nog gezet.

[	0.	0.	277714.	556593.	831557.	1111509.	1387420
	1666589.	1391111.	1109824.	834771.	556060.	276852.]	

Bewijst dit nu wél wat? Nee. Bewijzen doet zoiets zelden, zeiden we al. Maar als alle uitkomsten meer dan honderdduizend keer voorkomen en 0 en 1 niet, dan doet het wel iets vermoeden. Het bevestigt wat ik al dacht: dat je nooit minder dan 2 kunt gooien met twee dobbelstenen. Dat 11 twee keer zo vaak voorkomt als 12, begint er ook aardig op te lijken. (Deel de twee getallen op elkaar en je komt op 2,0087. Dat wijkt minder dan 1% af van 2.)

Maar het bewees toch allemaal niets? Nee, niet echt. Maar ik heb nu twee dingen gedaan.

1. goed nadenken over statistiek
2. mijn best doen op een goed programma dat het simuleert

Als wat daar uitkomt met elkaar overeenstemt, dan klopt het misschien wel. Meer weet je niet, maar het is al veel meer dan niets.

75. Eén dobbelsteen

Hoe groot is het aantal ogen van een dobbelsteen die je op tafel gooit? Voor een gewone dobbelsteen is het antwoord: 21. Die heeft namelijk zes kanten, met in totaal $1 + 2 + 3 + 4 + 5 + 6 = 21$ ogen.

Meestal bedoelen we met het aantal ogen niet het totale aantal ogen, maar het aantal ogen van de kant die boven ligt. Dan zijn er zes mogelijke antwoorden: 1, 2, 3, 4, 5 of 6. Als je het aantal ogen van één bepaalde worp bedoelt tenminste! Het aantal ogen *in het algemeen*, het aantal ogen van *een toevallige worp*, een worp die je *niet doet maar zou kunnen doen*; dat is een ander verhaal. Die vraag heeft maar één antwoord: 1, 2, 3, 4, 5 of 6, met elk van die uitkomsten een kans van $1/6$.

We hebben nu dus verschillende antwoorden op dezelfde vraag.

- 21: het totale aantal
- 5: het aantal ogen van de worp die ik zojuist deed
- 2: het aantal ogen van de worp die ik daarvoor deed
- 6: het aantal ogen dat ik hoopte te gooien
- $P(X)$: een uniforme discrete kansverdeling van 1 t/m 6

Veel mensen die statistiek lastig vinden, lijken moeite te hebben met het begrip *kansvariabele* of *toevalsvariabele*. Zo'n soort variabele heeft namelijk niet een bepaalde waarde, maar beschrijft alle mogelijke uitkomsten en de kans op elk van die mogelijke uitkomsten.

In de kansrekening is een stochastische variabele of stochastische grootheid een grootheid waarvan de waarde een reëel getal is dat afhangt van de toevallige uitkomst in een kansexperiment. De stochastische variabele, ook toevalsvariabele, kansvariabele of stochast, is een eigenschap van de uitkomst die in een getal is uit te drukken, zoals de leeftijd of het inkomen van een toevallige voorbijganger. Wikipedia

Er in dit soort definities altijd wel een mogelijkheid om ergens moeilijk over te doen. Een toevallige voorbijganger... bedoel je daarmee *die ene* toevallige voorbijganger, die vijf minuten geleden voorbijkwam? Die heeft heus wel één bepaalde leeftijd. Bijvoorbeeld 35 jaar. Of 21 jaar. Misschien wel 64 jaar of 11 jaar. Of 111 jaar, maar dat is minder waarschijnlijk. Hoewel... Als die ene voorbijganger toevallig 111 jaar was, dan is daar niets onwaarschijnlijk meer aan. Het is gewoon waar.

Het lastige van toevalsvariabelen is dat je hun uitkomst niet als een hard getal kunt berekenen. Het hoogst mogelijke aantal ogen met twee dobbelstenen

is makkelijk: dat is 12. Het gemiddelde aantal ogen met twee dobbelstenen is ook niet moeilijk: dat is 7. Maar dat zijn geen toevalsvariabelen, ook al is het gemiddelde van een verdeling wel iets wat in de statistiek wordt gebruikt.

De schatting van het gemiddelde op basis van een aantal worpen is wél een toevalsvariabele. Als je één keer gooit met één steen, komt er volstrekt willekeurig een 1, 2, 3, 4, 5 of 6 uit. Dat zegt maar heel weinig over het gemiddelde. Je zult nooit 0 of 7 gooien en ook geen -3 of 42. En geen 9,8 ook. Maar naarmate je vaker gooit, zal de uitkomst gaandeweg dichterbij de 3,5 liggen.

Wat gebeurt er met het gemiddelde aantal ogen als je na N keer gooien nog een keer gooit? Komt dat dichterbij het echte gemiddelde te liggen? Antwoord: dat weet je niet. Stel dat het gemiddelde tot nu toe 3,75 was. Dan komt het verder van het echte gemiddelde te liggen als 4, 5 of 6 gooit. En als je 1 gooit? Dan komt het er dichterbij. Hoewel... Als die 3,75 berekend is uit de vier worpen $6 + 2 + 3 + 3 = 14$, dan is dat juist niet zo. Want $14 / 4 = 3,75$. En $15 / 5 = 3$, dat verder van 3,5 ligt dan 3,75.

We zouden het programma kunnen uitbreiden door vaker dan één keer te werpen en de resultaten te bewaren in een array. Dan kunnen we niet alleen zien of alle waarden voorkomen, maar ook of dat even vaak gebeurt.

```
import random
import numpy as np

aantal_worpen = 10
uitkomst = np.zeros(7, dtype=int)

totaal = 0
for i in range(aantal_worpen):
    ogen = random.randint(1,6)
    uitkomst[ogen] = uitkomst[ogen] + 1

print(uitkomst)
```

Hé, een array met een lengte van zeven ... Er zijn toch maar zes mogelijke uitkomsten? Ja. Maar het is handig om het aantal gegooide ogen als index te nemen. Dan moet de index lopen van 0 t/m 6, omdat een array bij 0 begint.

In mijn eerste uitvoer van het programma kwam dit eruit.

```
[0 3 1 1 3 0 2]
```

Het getal 5 komt nooit voor. En 1 en 4 vaker dan de andere. Maar ja, een toevalsvariabele waar je maar tien keer naar kijkt terwijl er zes mogelijkheden zijn, dat is ook wel heel erg weinig. Laten we het honderd keer doen.

```
[ 0 12 12 17 21 15 23]
```

Het verschil tussen het hoogste aantal (23) en het laagste (12) is bijna een factor 2. Honderd is nog steeds niet erg veel. Maar waarom ook zo zuinigjes? Het kost nauwelijks tijd. Duizend of tienduizend keer... waarom niet gewoon een miljoen? En dan de uitkomst niet afdrukken als aantal, maar als percentage. Een geheel aantal procenten lijkt me wel mooi. Voor het laatste statement moeten we dan iets slims bedenken.

```
print(uitkomst * 100 // aantal_worpen)
```

De dubbele deelstreep om er een gehele deling van te maken. Als je hebt leren programmeren in C en daarna bent overgestapt op *Python*, is dat met die deelstreep altijd even opletten. In C is het resultaat van de enkele deelstreep afhankelijk van het type van de operanden. Als dat allebei integers zijn, komt er een integer uit. In *Python* is dat een float. Als je in *Python* een geheel getal wilt, moet je een dubbele deelstreep nemen. En als je een geheel aantal procenten wilt: eerst vermenigvuldigen met 100.

```
[ 0 16 16 16 16 16 16]
```

Dit is de uitkomst die ik verwachtte. Ja, er lijkt 4% te ontbreken. Maar dat komt door de afronding. Als je dat niet leuk vindt, moet je het laatste statement nog wat uitbreiden.

```
print((uitkomst * 1000 // aantal_worpen) / 10)
```

Dan wordt de uitvoer (althans in mijn eerste poging):

```
[ 0.  16.7 16.6 16.6 16.6 16.6 16.6]
```

De som is 99,7 en dat komt al wat dichterbij de 100. Het hele spel van meer herhalingen en grotere of kleinere aantallen decimalen kun je eindeloos blijven herhalen. Waar je tegenaan loopt, is dat een simulatie heel lang kan gaan duren. Of erger: dat de getallen te groot worden om nog opgeslagen te worden in een integer. Dat gebeurt niet zo heel snel, maar het houdt wel een keer op.

76. Betere schatting van spreiding

Als je een steekproef neemt uit een normaal verdeeld proces, dan kun je uit die steekproef een gemiddelde en variantie (en standaardafwijking) bepalen. We berekenen dan op grond van tien willekeurige getallen van het (in theorie) normaal verdeelde proces met gemiddelde μ en standaardafwijking σ . Dat gemiddelde van die ene steekproef duiden we aan met $\bar{x}$ en s_x .

Vaak doe je maar één steekproef en zorg je ervoor (of ga je er zonder nadenken van uit) dat die groot genoeg is om de waarden $\bar{x}$ en s_x dicht in de buurt te laten komen van wat je eigenlijk zou willen weten, namelijk μ en σ .

Wanneer zijn het er genoeg? Twee of drie is wel erg weinig en duizend klinkt als erg overdreven. Maar welke waarde daartussenin is goed genoeg?

Laten we voor de aardigheid eens tien verschillende steekproeven van tien gesimuleerde grootheden doen. De uitkomsten daarvan zetten we in een tabel. De kolommen zijn de steekproeven, de bovenste tien waarden zijn de gevonden metingen en onderaan staan de $\bar{x}$ en s_x per steekproef.

We zien in de tabel dat $\bar{x}$ varieert tussen 4,27 en 5,46 (een onderling verschil van 28 %) en dat s_x varieert tussen 0,73 en 1,21 (onderling verschil: 66 %).

4,59	5,55	3,76	4,56	6,54	5,60	5,91	5,65	4,76	4,37
4,43	6,74	6,94	5,21	3,75	5,87	4,96	5,77	5,77	4,34
3,44	6,30	4,92	5,85	5,23	5,86	3,57	4,64	5,82	4,20
3,59	5,48	6,01	3,67	5,23	5,47	4,26	3,86	5,27	4,76
3,15	6,80	3,43	6,79	6,02	6,27	3,93	3,41	5,14	3,53
5,74	4,98	5,18	3,25	4,52	7,04	5,16	4,87	4,09	4,79
5,01	3,89	5,41	3,95	5,81	3,56	7,31	6,29	4,91	5,15
4,14	4,58	5,64	5,31	5,44	5,90	6,41	5,97	4,88	5,38
5,32	5,09	5,41	3,50	4,57	4,11	6,15	4,56	8,37	4,00
3,27	4,53	4,77	6,02	5,33	4,88	5,96	5,32	5,37	2,98

$\bar{x}$	4,27	5,39	5,15	4,81	5,24	5,46	5,36	5,03	5,44	4,35
s_x	0,90	0,98	1,02	1,21	0,81	1,03	1,19	0,94	1,15	0,73

Het gemiddelde van die tien gemiddelden is overigens 5,05 en het gemiddelde van die tien standaardafwijkingen is 0,99. Dat wijkt allebei 1 % af van 5 en 1... Laten 5 en 1 nu precies de μ en σ zijn waarmee deze getallen gegenereerd zijn.

Als je diezelfde tien steekproeven nog een keer doet, zul je tien andere gemiddelden en tien andere standaardafwijkingen vinden. Het gemiddelde daarvan zal waarschijnlijk niet heel veel afwijken van 5 en 1.

Nu de wiskunde erachter. Wat hieronder volgt, is eigenlijk niet echt moeilijk, maar het is wat lastig om goed onder woorden te brengen.

Als x een stochastische variabele is die normaal verdeeld is met een gemiddelde van μ en standaardafwijking σ , dan zijn de $\bar{x}$ en s_x die uit een steekproef berekend worden zelf óók normaal verdeeld. De ene steekproef levert een ander gemiddelde op dan de andere en ook standaardafwijkingen verschillen. Het zijn zelf ook *toevalsvariabelen*, die afhangen van de toevallige uitkomst van die tien (of hoeveel dan ook) metingen.

Wat is het gemiddelde van al die gemiddelden? Dat weet je niet precies, maar de verwachtingswaarde is wel gelijk aan μ . De normale verdeling van die steekproefgemiddelden heeft namelijk dezelfde μ . Eigenlijk ook wel logisch, want de kans op een waarde die afwijkt naar boven is even groot als de kans op dezelfde afwijking naar beneden. Als je oneindig vaak zou middelen, is de gemiddelde waarde van $\bar{x}$ – net als die van x – gelijk aan μ .

Hoe groot is de *standaardafwijking* van dat gemiddelde? Die vraag is wat minder simpel, maar gevoelsmatig zeg je dat die kleiner zal zijn dan de standaardafwijking van de verdeling zelf. In een gemiddelde van tien metingen zal minder spreiding zitten dan in de meting zelf.

Als je de lengte van studenten in één klas op een rij zet, zul je een flinke spreiding zijn. Afhankelijk van wat je 'flink' noemt. Maar vaak zit er wel iemand van 2 meter of langer tussen en je hebt ook studenten van 1,60 meter of kleiner. Als je de gemiddelde lengte van alle studenten in één klas vergelijkt met de gemiddelde lengte van alle studenten in andere klassen, zal die spreiding veel minder groot zijn. Je vindt nooit een klas waar studenten gemiddeld 2 meter lang zijn, tenzij je ze heel bewust geselecteerd hebt. De spreiding van de gemiddelden is nu eenmaal kleiner dan de spreiding van de individuele waarden.

Wiskundig gezien kun je bewijzen dat de standaardafwijking van het gemiddelde gelijk is aan.

$$\sigma_{\bar{x}} = \frac{\sigma_x}{\sqrt{N}}$$

De formule is wat groot afgedrukt, omdat hij nogal belangrijk is. Daarom is het ook zo belangrijk dat je snapt wat er staat.

Om met de makkelijkste te beginnen: N is de grootte van de steekproef. Dan die ene die we al wat langer kennen: σ_x , de standaardafwijking van alle mogelijke metingen die je kunt doen. (Van de onderliggende verdeling dus.)

De lastigste is $\sigma_{\bar{x}}$, die met dat streepje boven de x . Dat is namelijk de *standaardafwijking van het steekproefgemiddelde*, in het Engels vaak aangeduid als *Standard Deviation of the Mean* of kortweg *Standard Error*.

Zoals σ_x een handige maat is om aan te geven hoe groot de spreiding is van de *individuele metingen*, zo is $\sigma_{\bar{x}}$ een handige maat voor de spreiding in de *gemiddelden van de steekproeven*. Het is meestal een getal dat je lekker klein wilt krijgen. Want als de spreiding in de gemiddelden van de steekproeven klein is, dan maakt het niet zo veel meer uit welke steekproef je neemt. Er komt toch telkens ongeveer hetzelfde uit.

Die σ_x die in de teller van de breuk staat, daar doe je meestal niet zo veel aan. Dat is gewoon de spreiding die in je metingen zit. Die N is een ander verhaal. Vaak kun je zelf kiezen hoe groot je steekproef is.

Het is belangrijk om in te zien dat die N een getal is waarvan eerst de wortel wordt getrokken voordat er gedeeld wordt. Als je zou willen dat de spreiding twee keer zo klein wordt, zul je dus vier keer zoveel metingen moeten doen.

77. De spreiding in de spreiding

In het vorige hoofdstuk zagen we dat het gemiddelde van een steekproef van normaal verdeelde metingen zelf ook normaal verdeeld is. Het gemiddelde van die verdeling is μ . Als de gemiddelde lengte in een groep vrouwen 1,70 m is, dan zal het gemiddelde van een willekeurige steekproef ook 1,70 m zijn.

De spreiding van de steekproefgemiddelden is kleiner en kan worden berekend door te delen door de wortel van de steekproefgrootte.

$$\sigma_{\bar{x}} = \frac{\sigma_x}{\sqrt{N}}$$

Als de standaardafwijking in de lengte van die groep vrouwen 5 cm is, dan zal de standaardafwijking in het gemiddelde van een steekproef van 25 gemiddeld genomen maar 1 cm zijn.

Uiteraard geldt ook hier: de ene groep van 25 vrouwen is de andere niet.

De uitspraak is natuurlijk niet te ‘bewijzen’ met een simulatie, maar het is niet heel moeilijk om (bijvoorbeeld in *Excel*) een tabel te genereren met tien ‘steekproeven’ onder 25 fictieve vrouwen die een normaal verdeelde lengte hebben met een gemiddelde van 1,70 m en een standaardafwijking van 5 cm. De resultaten van de gemiddelden en standaardafwijkingen per steekproef zijn als volgt.

171	171	170	169	169	171	169	172	170	169
5,1	6,0	5,1	5,6	4,7	3,8	7,0	6,3	5,5	3,9

Natuurlijk kun je uit deze twee regels weer een gemiddelde bepalen: het gemiddelde van de steekproefgemiddelden is 1,702 m en de standaardafwijking van de steekproefgemiddelden is 0,90 m. (We verwachtten 170,0 en 1,0 m. Het komt in de buurt. En het bewijst niets, maar geeft een goed gevoel en misschien wat meer begrip.)

Nu de laatste stap, dus hou nog even vol. Niet alleen het steekproefgemiddelde verschilt per steekproef, maar ook de spreiding (standaardafwijking) is voor elke steekproef anders. En omdat die voor elke steekproef anders is, is er sprake van een spreiding in de standaardafwijkingen. Eigenlijk zou je die spreiding ook willen kennen. Alleen als de spreiding in de spreidingen klein is, heb je er voldoende aan om er eentje te bepalen.

Ben je er nog? Mooi. Hier komt de formule.

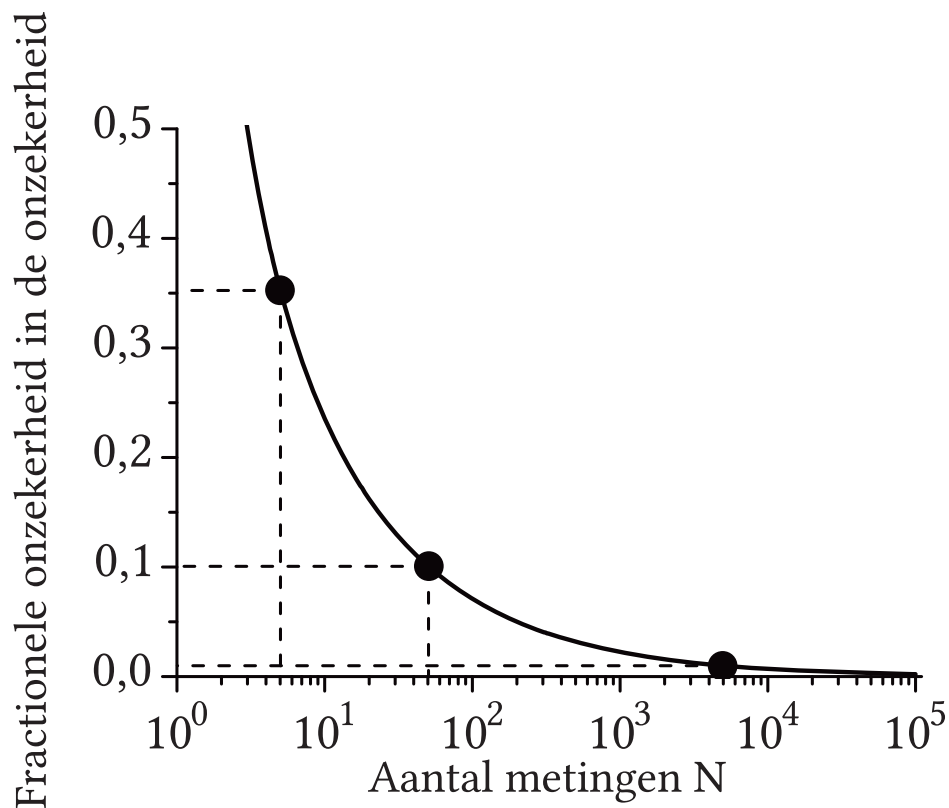
$$\sigma_s = \frac{\sigma}{\sqrt{2N-2}}$$

In ons geval was σ gelijk aan 5 cm en N gelijk aan 25. Invullen:

$$\sigma_s = \frac{\sigma}{\sqrt{2N-2}} = \frac{5}{\sqrt{2 \cdot 25 - 2}} \approx 0,72 \text{ cm}$$

Wat het een beetje ingewikkeld maakt: we kennen wel de N in de formule, maar in praktijk meestal niet de σ . Als we niet wisten dat hij 5 cm was (wat we wel weten, we hebben immers zelf die simulatie geschreven), dan zouden we schatten dat hij 5,3 cm was, omdat dat het gemiddelde is van de standaardafwijkingen die we uit de steekproeven haalden. Die 0,72 cm helpt ons om een beeld te vormen hoe betrouwbaar die 5,3 cm was.

Hughes & Hase nemen in hun boek een grafiekje op dat een indruk geeft van het verband tussen het aantal metingen en de (fractionele) onzekerheid in de onzekerheid. Zij hebben het dan over de *fractional error in the error*, maar omdat ze het over de spreiding (standaardafwijking) hebben, bedoelen ze niet de *fout*, maar de *onzekerheid*.



Figuur 77.1 De afname van de onzekerheid bij een toename van het aantal metingen

Getallen in een tabel kunnen ook helpen om een gevoel te krijgen.

N	$\frac{1}{\sqrt{N}}$	$\frac{1}{\sqrt{2N-2}}$
1	100%	
2	71%	71%
3	58%	50%
4	50%	41%
5	45%	35%
10	32%	24%
20	22%	16%
50	14%	10%
100	10%	7%
500	4%	3%
1000	3%	2%
10000	1%	1%

Wie één meting doet, heeft een spreiding die net zo groot is als de spreiding in de verdeling zelf. Door vier metingen te doen en te middelen, halveert de spreiding. Je hebt er honderd nodig om tot 10% te komen en tienduizend voor 1%.

In de rechter kolom staat aangegeven hoe goed je die standaardafwijking dan kent op grond van de metingen. Bij vijftig metingen kun je zeggen dat je zo'n 10% onnauwkeurigheid hebt in het schatten van je spreiding. Dat is vaak goed genoeg.

78. Rough-and-ready approach

Hughes & Hase beschrijven in hun boek een alternatief voor het berekenen van de standaardafwijking, in de vorm van een grove schatting die ze de *rough-and-ready* approach noemen. *Rough-and-ready* is een Engelse aanduiding die niet echt goed in het Nederlands te vertalen is. Vaak wordt het vertaald met ‘basaal’, ‘ruw’ of ‘pragmatisch’. Het is een ruwe schatting (*rough*) en je bent lekker snel klaar (*ready*).

Ze leggen hun methode uit aan de hand van een voorbeeld, waarin er tien slingertijden worden gemeten.

nr.	1	2	3	4	5	6	7	8	9	10
T (s)	10,0	9,4	9,8	9,6	10,5	9,8	10,3	10,2	10,4	9,3

De steekproefstandaardafwijking hiervan vind je vrij eenvoudig door ze te kopiëren en plakken naar *Excel* en de formule =STDEV.S(A1:A10) toe te passen. De waarde die je dan vindt, is 0,419. In het vorige hoofdstuk hebben we gezien dat de relatieve fout in een schatting met tien waarden 24% is. De standaardafwijking zal dus liggen tussen 0,32 en 0,52.

De *rough-and-ready approach* is handig als je geen *Excel* bij de hand hebt of niet weet hoe je een standaardafwijking berekent uit een steekproef. Hij werkt tamelijk eenvoudig.

Het gemiddelde van de metingen is 9,93 s, afgerond op één decimaal dus 9,9 s. Dat betekent dat de hoogste waarde 0,6 s hoger is dan het gemiddelde en de laagste waarde 0,6 s lager. Je krijgt het vermoeden dat Hughes & Hase opzettelijk waarden hebben gekozen waarbij de uitersten keurig symmetrisch zijn. En ook nog eens op een lekker deelbare 6 uitkomen. Bij de *rough-and-ready approach* kijk je namelijk naar de grootste afwijking van het gemiddelde. Die is in dit geval in beide ‘richtingen’ gelijk aan 0,6 s. Daar neem je 2/3 van als schatting voor de standaardafwijking. Dat is dus 0,4 s.

Als Hughes & Hase met een decimaal meer hadden gerekend, was het sommetje lastiger uit het hoofd te doen. Dan waren de grootste afwijkingen niet 0,6 s en 0,6 s geweest, maar 0,63 s en 0,57 s. Dan hadden ze 2/3 van 0,63 s moeten nemen en was het antwoord 0,42 s geweest. Maar waarom zou je twee significante cijfers nemen als het zo grof is en je een onzekerheid van een kwart hebt?

In formulevorm zijn dit de stappen die je neemt bij de *rough-and-ready approach*.

Bereken het gemiddelde van de data set, $\bar{x}$.

Bereken de maximale afwijking van dat gemiddelde. Dat kan dus $d_{max} = x_{max} - \bar{x}$ zijn of $d_{max} = \bar{x} - x_{min}$, afhankelijk van welke twee het grootste is.

Als standaardafwijking neem je dan $\sigma = 2/3 \times d_{max}$

Een beginner vindt het misschien vreemd dat je $2/3$ neemt van de maximale afwijking. Waarom niet gewoon de maximale afwijking zelf? Als je zegt dat de slingertijd gelijk is aan $9,9 \pm 0,6$ s, dan vallen alle gemeten waarden precies binnen dat bereik.

Op zichzelf is dat waar, maar het is gebruikelijk om de (geschatte) *standaardafwijking* van een normaal verdeeld proces op te geven als aanduiding voor de onzekerheid. Met die onzekerheid geef je niet aan dat alle waarden binnen dat bereik liggen, maar de meeste waarden wel. Bij een normale verdeling ligt 68% van de waarden binnen één standaardafwijking van het gemiddelde.

Hughes & Hase zeggen hierover:

The factor of two thirds is justified for a Gaussian distribution, discussed in Chapter 3. Note that the spread of the measurements is significantly worse than the precision of the measuring instrument – this is why taking repeat measurements is important.

Dat laatste is zeker waar. Denk nou niet dat de onzekerheid van je meetinstrument het (enige) probleem is. Meestal zit er ook nog een belangrijke bijdrage in degene die de meting doet en datgene wat je meet.

Maar nu dat eerste. Wordt een factor twee derde inderdaad gerechtvaardigd door de gaussverdeling? Het is waar dat $2/3$ ($\approx 0,6667$) niet veel afwijkt van $0,6827$. Dat verschil is $0,016$ oftewel $2,4\%$. Maar vergelijken we nu geen appels met peren?

Die 68% van de normale verdeling slaat op het percentage van de metingen dat minder dan één standaardafwijking van het gemiddelde ligt.

Die $2/3$ van de *rough-and-ready approach* slaat op het deel van grootste afwijking van het gemiddelde. Wat is bij een normale verdeling de grootste afwijking van het gemiddelde? Oeps... Die bestaat niet. Een normale verdeling loopt oneindig ver door.

Vaak zeggen we dat 68% binnen $\pm\sigma$ ligt en 95% tussen $\pm 2\sigma$. Je leest ook wel dat 99% tussen $\pm 3\sigma$ ligt, maar verder hoor je zelden iemand. Eigenlijk is die 99% wat grof (en zuinig) uitgedrukt en zou je $99,73\%$ moeten zeggen.

We kunnen de rij nog langer maken. Hieronder staat hoeveel procent van de data (kolom 1) binnen N (kolom 2) keer de standaardafwijking ligt. Kolom 3 vertelt dat één op de zoveel erbuiten valt.

68%	1	3
95,5%	2	22
99,73%	3	370
99,9937%	4	16 000
99,999943%	5	1,7 miljoen
99,99999980%	6	een half miljard
99,9999999974%	7	een half biljoen

Voor acht en meer keer de standaardafwijking lukt het *Excel* niet meer om het uit te rekenen. Maar ze komen theoretisch wel voor.

De *rough-and-ready approach* is dus wel heel grof in zijn benadering en heel grof in zijn uitleg. Een alternatief zou zijn dat je bij tien metingen de grootste drie afwijkingen buiten beschouwing laat.

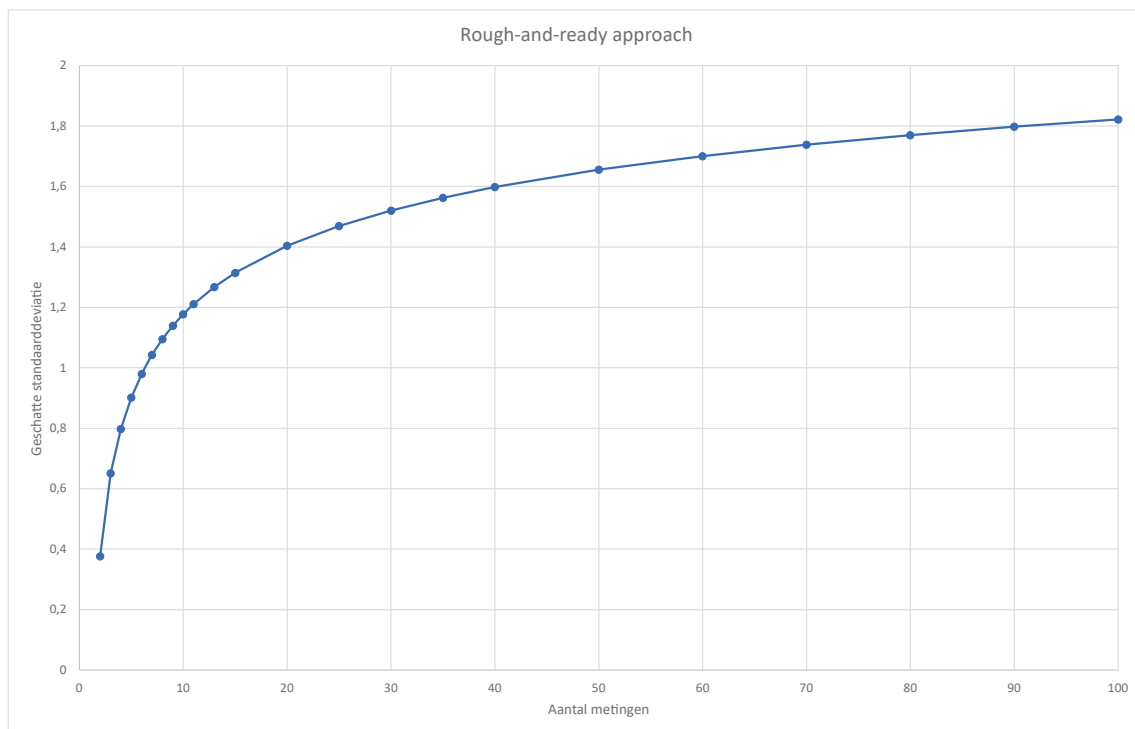
Laten we het tabelletje hierboven eens sorteren van klein naar groot.

nr.	10	2	4	3	6	1	8	7	9	5
T (s)	9,3	9,4	9,6	9,8	9,8	10,0	10,2	10,3	10,4	10,5
Δ	0,63	0,53	0,33	0,13	0,13	0,07	0,27	0,37	0,47	0,57

Als grootste drie afwijkingen weglaat, liggen de andere zes tussen 9,6 en 10,4. Het verschil tussen die twee zou je als twee keer de standaardafwijking kunnen beschouwen. Je komt dan nog steeds uit op 0,4 s.

We zouden met behulp van simulaties kunnen kijken hoe goed de schattingen met *rough-and-ready* eigenlijk zijn en het verband met het aantal metingen kunnen bepalen. Waar veel schatters beter worden naarmate je meer metingen hebt, is dat bij deze methode bepaald niet vanzelfsprekend. De grootste afwijking van het gemiddelde zal namelijk als maar groter worden naarmate je meer metingen hebt. Natuurlijk is de methode daar niet voor bedoeld, maar als je miljoenen of miljarden metingen hebt, komen er vast wel een waarden van vijf of zes keer de standaardafwijking voor.

Als je een simulatie een paar miljoen keer uitvoert op setjes van telkens N standaardnormaal verdeelde meetresultaten, dan zie dat er altijd in drie decimalen hetzelfde getal uitkomt als schatting voor de standaardafwijking. Bij kleine steekproeven is die namelijk te klein: de maximale waarde van een kleine steekproef heeft nu eenmaal een kleinere verwachtingswaarde dan die van een grote steekproef. De grafiek die resulteert, vertoont een continu stijgende lijn.



De *rough-and-ready* geeft gemiddelde genomen aardige resultaten voor vijf tot acht metingen, maar is bij een kleiner aantal gemiddeld te laag en bij een groter aantal gemiddeld te hoog. Dat is op zich niet heel erg, want op aantallen die kleiner zijn is het gewoon niet zo zinnig om een schatting te doen, en als je groter steekproeven neemt is er niet echt een goede reden meer om een schatting te doen.

De grafiek blijft maar stijgen en zal zelfs boven de 2 uitkomen. Dat gebeurt namelijk als de grootste waarde 3 is en dat is bij een steekproef van 1000 in ongeveer driekwart van de gevallen zo.

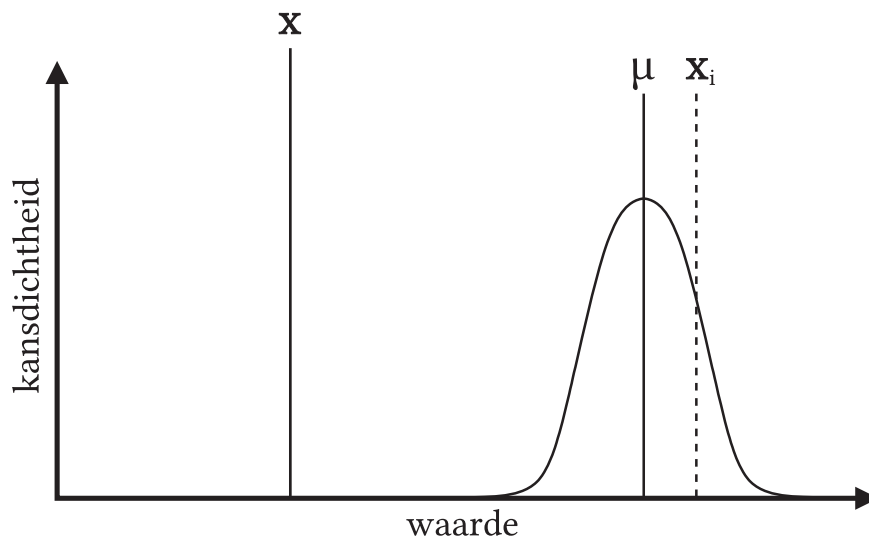
79. Kansverdeling resultaat en meting

Kansvariabelen zijn lastige dingen. Gemiddelden ook trouwens. Bedoel je het “echte” gemiddelde van de verdeling, of dat ene “toevallige” gemiddelde van die ene steekproef?

Dan heb je ook nog *standaardafwijkingen*. De kansverdeling heeft een standaardafwijking, die je de “echte” zou kunnen noemen. De steekproef heeft er ook eentje. En net zoals je per steekproef een ander gemiddelde vindt, vind je per steekproef ook een andere standaardafwijking.

Het woord “echt” kan verwarring geven. Die kansverdeling is niet iets wat werkelijk “bestaat”. In die zin is het gemiddelde van de kansverdeling helemaal niet *echt*, maar *theoretisch* en verzonnen. Het is een *model* van de werkelijkheid. In de werkelijkheid hebben we wel concrete metingen, maar geen abstracte kansverdelingen. Die heb je alleen in de wiskunde en de statistiek.

Meting en *meetresultaat* moet je ook nog uit elkaar houden, al is het verschil tussen die twee heel anders. Er is geen “echte” meting tegenover een “echt” meetresultaat. Het verschil tussen die twee zit in de aantallen. Een meting is er altijd één. Een meetresultaat is het gemiddelde van een aantal metingen.



Figuur 79.1 Kansverdeling van één meting

De waarde x in Figuur 79.1 is de werkelijke waarde (*true value*). Bestaat de werkelijke waarde ook echt? Ja, maar we kennen hem meestal niet. (Anders hoefden we niet te meten.) Wat we soms wel kennen, is de geaccepteerde waarde. Die zal dicht in de buurt van x liggen, maar niet exact hetzelfde zijn. De werkelijke waarde x (die we niet kennen) heeft geen onzekerheid.

De geaccepteerde waarde (die niet in de grafiek is getekend) heeft wel een onzekerheid. Die zal meestal veel kleiner zijn dan die van de meting.

Accepted values heb je van natuurconstanten, maar niet van een houten plank die je zelf op een zaterdagmiddag op lengte hebt gezaagd. Daar bestaan ook geen handleidingen van waarin je die lengte (of de onzekerheid erop) zou kunnen opzoeken. Je weet van de lengte van die plank helemaal niets, totdat je gemeten hebt. Nou ja, je kunt natuurlijk wel wat schatten. Maar zonder te meten sla je zomaar de plank mis.

We kennen x dus niet. Wat we wel kennen, is de uitkomst van één bepaalde meting x_i . Die ene x_i beschouwen we als de toevallige uitkomst van een kansvariabele. Hij had ook een andere waarde kunnen hebben. Sterker nog: als je opnieuw een meting doet, zal die zeer waarschijnlijk anders zijn, als er tenminste sprake is van een toevallige fout.

De kromme in Figuur 79.1 is de kansdichtheidsfunctie van de meting. Dat is een niet “echt bestaande”, maar theoretische wiskundige beschrijving (modellering) van de werkelijkheid van het meten.

In Figuur 79.1 is die verdeling normaal, maar dat hoeft niet zo te zijn. De verdeling kan ook uniform zijn. Of één bepaalde waarde hebben, als er geen toevallige fout is. Dan is de uitkomst van de meting altijd hetzelfde, al herhaal je nog zo vaak. Maar nog steeds is de meting x_i dan niet de echte x .

Het is een veelgemaakte denkfout: “Als er geen spreiding is, dan is er geen onzekerheid...” Die is er wél dus! Je kunt die alleen niet schatten door te herhalen.

Als je één enkele meting aan een grootte doet en je hebt geen aanvullende gegevens (uit specificaties of kalibraties), dan heb je niets anders dan één getal en een eenheid.

We weten niet hoe ver x_i van μ vandaan ligt.

We weten niet hoe ver μ van x vandaan ligt.

We weten niet hoe groot σ is.

We hebben alleen een meting x_i . Die heeft een waarde (die we bepaald hebben) en een onzekerheid (waar we geen flauw idee van hebben.)

Nu gaan we de stap maken van een *meting* naar een *meetresultaat*. We doen nog negen metingen en hebben dan tien uitkomsten. Die kunnen we middelen en we kunnen de standaardafwijking van de steekproef bepalen. Het gemiddelde vind je door gewoon alles op te tellen en te delen door (in dit geval)

tien. De standaardafwijking van de steekproef vind je met de formule die we daarvoor kennen.

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

In deze formule staan geen μ en σ . Die kennen we namelijk niet. Er staan wel een s (die we berekenen) en een $\bar{x}$ (die we berekend hebben).

Alle symbolen nog even op een rij.

μ : het gemiddelde van de (theoretische) verdeling

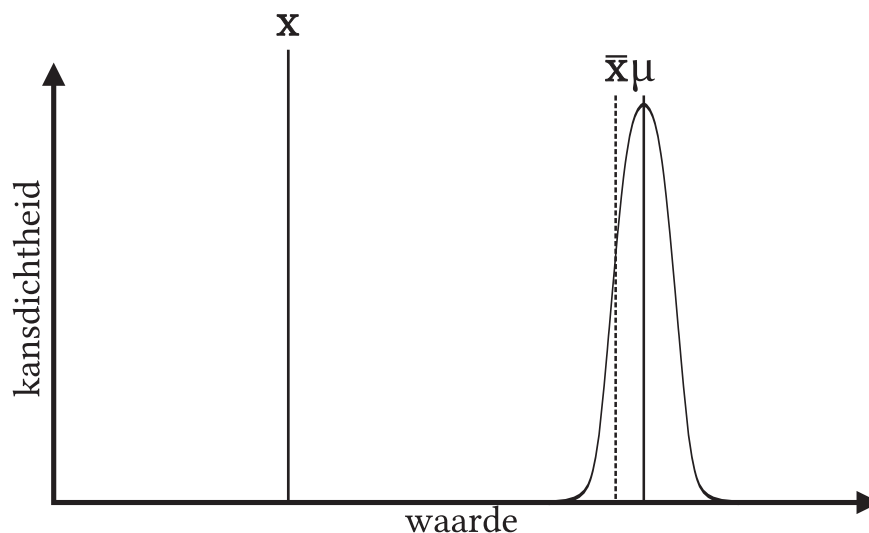
σ : de standaardafwijking van de (theoretische) verdeling

$\bar{x}$: het gemiddelde van de (praktische) steekproef

s : de standaardafwijking van de (praktische) steekproef

Weten we nu meer? Zijn we iets opgeschoten door de meting tien keer uit te voeren? Jazeker! Een belangrijke winst is dat we met s een mooie schatting in handen hebben van σ , en dus van de grootte van de toevallige fout. Een andere belangrijke winst is dat $\bar{x}$ naar verwachting dichterbij de echte μ zal liggen dan die ene x_i die we hadden.

Dat brengt ons op Figuur 79.2.



Figuur 79.2 Kansdichtheidsfunctie van een meetresultaat (gemiddelde van een aantal metingen)

Zoek de verschillen!

- De normale verdeling is veel smaller (en hoger).
- Er staat geen x_i in het plaatje, maar een $\bar{x}$.

De normale verdeling in deze afbeelding is niet de kansverdeling van *een meting*, maar van *een meetresultaat van tien metingen*.

Het gemiddelde van (in dit geval) tien normaal verdeelde metingen is zelf ook normaal verdeeld. De standaardafwijking ervan is een stuk kleiner. Die is te berekenen met de formule die we eerder hebben gezien.

$$\sigma_{\bar{x}} = \frac{\sigma}{\sqrt{N}}$$

Kennen we nu de standaardafwijking van de kansverdeling? Nee. Kennen we het gemiddelde? Ook niet. Maar het gemiddelde uit de steekproef van tien zal naar verwachting wel aardig in de buurt komen van het echte gemiddelde. En de standaardafwijking van die steekproef van tien geeft een aardig beeld van de standaardafwijking van de verdeling.

Hoe goed is het beeld dat we hadden van de standaardafwijking van de verdeling toen we één meting hadden gedaan? Antwoord: we hadden geen flauw idee! Eén meting heeft geen spreiding, want er zijn geen verschillen.

Hoe goed kennen we nu het gemiddelde van de verdeling? Het is de uitkomst van een normaal verdeeld proces met μ en σ . Die kennen we niet. We kennen wel s en $\bar{x}$. Dat zijn de beste schattingen die we hebben, en als het aantal metingen groot genoeg is, zijn die niet eens zo slecht. De onzekerheid op het gemiddelde schatten we dus gelijk aan s .

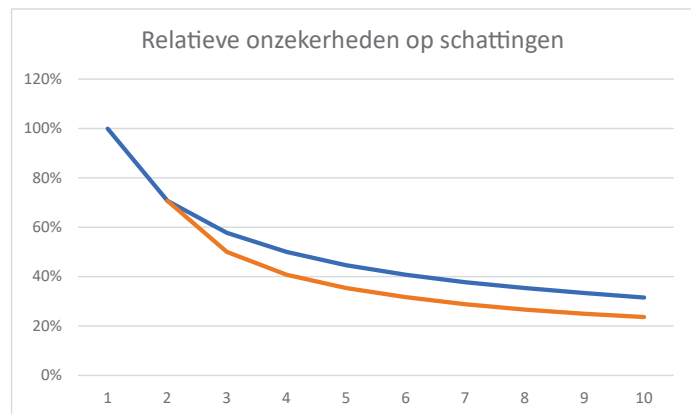
Hoe goed kennen we nu de standaardafwijking van de verdeling? Volgens de wiskunde volgt de relatieve onzekerheid hierin uit onderstaande formule.

$$\frac{1}{\sqrt{2N-2}}$$

Als N dus “aardig groot” is, wordt deze term klein. Dan kunnen we de onzekerheid met een kleine onzekerheid vaststellen.

Om een indruk te krijgen van wat die formules in praktijk voorstellen, zetten we wat waarden in een tabel en in een grafiek. Merk op dat in de tweede formule de waarde voor $N = 1$ niet gedefinieerd is.

	$\frac{1}{\sqrt{N}}$	$\frac{1}{\sqrt{2N-2}}$
1	100%	
2	71%	71%
3	58%	50%
4	50%	41%
5	45%	35%
6	41%	32%
7	38%	29%
8	35%	27%
9	33%	25%
10	32%	24%



Om het gemiddelde en standaardafwijking goed te schatten, is het doen van tien metingen nog wat weinig. Je zit er dan zomaar een derde of een kwart naast (wat soms overigens goed genoeg is). Voor een serieuzere statistische analyse moet je eerder denken aan honderd metingen.

Oefeningen

1. Hoeveel metingen moet je doen om het gemiddelde met een onzekerheid van maximaal 5% te kunnen schatten?
2. Hoeveel metingen moet je doen om de standaardafwijking met een onzekerheid van maximaal 10% te kunnen schatten?
3. Welke waarde ligt (in het algemeen) dicht bij het gemiddelde van de verdeling: x_i of $\bar{x}$?

80. Overschrijdingskans en significantie

Als een verdeling echt normaal is, zou je theoretisch ook 3, 4, 5 of 6 keer de standaardafwijking kunnen afwijken van het gemiddelde. Laten we voor het gemak eens stellen dat de basketballers uit het vorige hoofdstuk gemiddeld 2,00 meter lang zijn, met een standaardafwijking van 5 cm, is er dan ook een kans dat iemand 4,05 meter lang is? Of een negatieve lengte heeft: -5 cm... Het antwoord is: ja. Er bestaat een kans dat je 41 keer de standaardafwijking afwijkt van het gemiddelde. Alleen is die extreem klein.

Je kunt dat afdrukken in een tabel. In de eerste kolom zetten we een lengte in meters. In de tweede kolom het aantal keren de standaardafwijking boven het gemiddelde. In de derde, vierde en vijfde kolom noteren we de kans dat iemand minstens die lengte heeft. In elke van die drie kolommen staat hetzelfde getal, maar telkens in een andere notatie.

- De fractie, uitgedrukt in 20 decimalen
- Het percentage, met drie significante cijfers
- De fractie, in wetenschappelijke notatie

l [cm]	N	P	P	P
200	0	0,50000000000000000000	50,0%	$5,0 \times 10^{-1}$
202,5	0,5	0,30853753872598800000	30,9%	$3,1 \times 10^{-1}$
205	1	0,15865525393145800000	15,9%	$1,6 \times 10^{-1}$
207,5	1,5	0,06680720126885750000	6,68%	$6,7 \times 10^{-2}$
210	2	0,02275013194817910000	2,23%	$2,3 \times 10^{-2}$
212,5	2,5	0,00620966532577616000	0,621%	$6,2 \times 10^{-3}$
215	3	0,00134989803163010000	0,135%	$1,3 \times 10^{-3}$
217,5	3,5	0,00023262907903554000	0,023%	$2,3 \times 10^{-4}$
220	4	0,00003167124183312000	0,003%	$3,2 \times 10^{-5}$
222,5	4,5	0,00000339767312473871	0,0%	$3,4 \times 10^{-6}$

225	5	0,00000028665157192354	0,0%	$2,9 \times 10^{-7}$
227,5	5,5	0,00000001898956247803	0,0%	$1,9 \times 10^{-8}$
230	6	0,00000000098658770042	0,0%	$9,9 \times 10^{-10}$
232,5	6,5	0,00000000004015998645	0,0%	$4,0 \times 10^{-11}$
235	7	0,00000000000127986510	0,0%	$1,3 \times 10^{-12}$
237,5	7,5	0,00000000000003186340	0,0%	$3,2 \times 10^{-14}$
240	8	0,00000000000000000000	0,0%	$0,0 \times 10^0$

Huh? Die laatste uitkomst van $0,0 \times 10^0$ is toch zeker een tikfout in het boek? Nee. Dan hadden we die fout er namelijk wel uit gehaald in plaats van een grijze alinea eraan te besteden. Het is de uitkomst die *Excel* berekent. Kennelijk lukt het niet meer met zulke kleine getallen.

Deze tabel is een leuke oefening in het kiezen van notaties. Wat is een handige weergave? Wat geeft snel overzicht en hoeveel cijfers gebruik je?

In de kolom met de fracties gebeurt iets geeks. Bij 200 t/m 220 eindigt de waarde steeds met minimaal drie nullen. Daarna lijken er opeens cijfers bij te komen. Microsoft beweert dat *Excel* voor opslaan en berekenen 15 significante cijfers gebruikt. Dat is in overeenstemming met wat we bij die eerste getallen zien. Maar het geeft nog geen verklaring voor die ‘sprong’ met drie decimalen. Het lijkt wel over er ergens in de berekening iets gebeurt dat we niet overzien.

Percentages zijn vaak handig om een snelle indruk te geven. Dat gaat pas mis als de getallen zo klein worden dat je niet meer met maximaal drie cijfers achter de komma kunt weergeven hoeveel het is. Voor wetenschappelijke notatie geldt het omgekeerde. Een fractie $5,0 \times 10^{-1}$ is wel erg cryptisch voor 50% of een kans van $\frac{1}{2}$. En $6,2 \times 10^{-3}$ leest ook lastiger dan 0,62% of 0,6% of 6‰. Vooral in mondelinge communicatie: ‘Die kans is maar zes promille.’ Misschien zou je zelfs liever ‘een half procent’ zeggen, of simpelweg ‘minder dan één procent’.

Eigenlijk doen we zoiets in de statistiek altijd al wanneer we het over ‘plus of min twee sigma’ hebben. De kans dat een waarde niet meer dan twee keer de standaardafwijking verschilt van het gemiddelde wordt vaak 95% genoemd

en de kans dat hij meer verschilt is dan 5% en misschien wel verwaarloosbaar. Eigenlijk is het geen 5% maar (in tien decimalen) 4,5500263896%.

De wetenschappelijke notatie is erg handig voor alle getallen in de onderste helft van de tabel, waar de numerieke waarden zo klein worden dat het lastig is om uit te spreken of je een voorstelling te maken. Een kans van 0,00000001898956247803 zegt niet zoveel, ook al staan er een heleboel cijfers. En 0,0% zegt al helemaal weinig. Het voordeel van $1,9 \times 10^{-8}$ is dat alleen al die macht voldoende is. Twee significante cijfers zijn vrij standaard. Eentje is al voldoende om aan te geven hoe groot het getal ongeveer is, en met twee heb je dat al aardig scherp vastgelegd. Voor minder wetenschappelijk lezerspubliek is het beter om in plaats van 2×10^{-8} te schrijven dat de kans 1 op 50 000 000 is.

Tot slot: *Excel* bleek bij het berekenen van bovenstaande tabel het niet meer aan te kunnen. Uit de laatste weergegeven waarde komt domweg 0. Dat is niet vreemd, maar wat zegt het over de betrouwbaarheid van het getal ervoor?

81. Verwaarloosbare kans

Wanneer is een kans verwaarloosbaar? Eigenlijk nooit. Althans: er is geen algemeen getal voor te geven, want het hangt af van *het belang* of iets gebeurt.

Als je speling neemt om op tijd in de les te komen, hou je rekening met de gevolgen als het niet lukt. Te laat komen is hinderlijk voor je medestudenten, maar als je héél zachtjes op je sokken binnensluipt en blozend op de stoel meteen naast de deur gaat zitten, valt het met de consequenties wel mee. Te laat zijn voor een tentamen heeft tot gevolg dat je niet mee mag doen.

Hoeveel speling neem je om de bus te halen? Nooit zoveel dat de kans 100,00% is dat je hem haalt. Want als werkelijk alles tegenzit – er staan zes huizen in brand en daardoor zijn bijna alle straten in de wijk afgezet – doe je er heel veel langer over. Mis je de bus, dan neem je gewoon de volgende. Die komt over een kwartier of een half uur. Of over zesenhalf uur, als het de laatste bus van die dag was. Dan neem je wat meer speling en minder risico.

Als 5% van de mensen die een nieuw type pijnstiller slikt eraan overlijdt, kun je beter paracetamol blijven gebruiken. (Je hoofdpijn is trouwens wel weg als je overlijdt.) Zelfs als het 1% of 1‰ is. Maar als 5% van de patiënten milde bijwerkingen en klachten heeft terwijl het medicijn de kans op genezing van een ernstige ziekte halveert, is er alle reden om het wél te gebruiken.

Oefeningen

Je bent voor de Vierdaagse van Nijmegen verantwoordelijk voor de veiligheid van de wandelaars. In het parcours ontdek je een ontbrekende tegel die slecht zichtbaar is. Je schat de kans dat een deelnemer erdoor ten val komt en iets breekt erg klein. De kans dat iemand dit op tijd ziet, schat je op 99,99%.

1. Bereken de kans dat alle 47 000 wandelaars niet ten val komen door die ontbrekende tegel.
2. In de eerste editie van *Met en & Weten* stond een fout in dit hoofdstuk. Er was toen sprake van 47 000 000 wandelaars... Hoe groot zou de kans bij dat aantal zijn?
3. Als de kans exact 99,99% zou zijn, hoeveel wandelaars mogen er dan meedoen om de kans kleiner dan 99% (of 1%) te laten zijn dat er iemand valt over die ontbrekende tegel?

Gek eigenlijk... Je kunt toch niet vallen over een ontbrekende tegel? Nee. Je valt waarschijnlijk over de volgende tegel (die *niet* ontbreekt).

Toevallige fouten

82. Systematische fout door toeval (1)

Als een *toevallige* fout in een meting veroorzaakt wordt door een *systematische* afronding omlaag (lees: door het weglaten van decimalen), dan bevat deze fout een systematische component. Dat komt doordat het gemiddelde van het stuk dat je weglaat niet nul is.

Dit is er eentje om in de gaten te houden. We kennen de kansverdeling van de fouten in werkelijkheid eigenlijk nooit, maar we maken er wel schattingen van. Dat resulteert in normale of uniforme of andere verdelingsdichtheidsfuncties die altijd een gemiddelde van nul hebben. Het zou ook een beetje raar zijn als dat niet zo was. “Ik denk dat de fout die we maken gemiddeld -3 mV is...” Als je dat denkt, waarom trek je die waarde dan niet overal van af? De spreiding verandert er niet door en de gemiddelde fout wordt (in absolute zin) kleiner. Kleiner kan hij niet zijn: exact 0 mV.

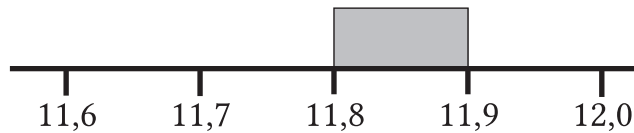
Stel dat je een multimeter hebt die de cijfers die hij niet kan weergeven domweg weglaat. Wat betekent het dan als je $11,8$ V meet? Ervan uitgaande dat de fout alleen maar wordt veroorzaakt door het weglaten van cijfers (wat in praktijk natuurlijk niet zo is) weet je dat de spanning minstens $11\,800$ mV is en maximaal $11\,899$ mV. (Met achter de komma nog een heleboel negens.)

De waarde die weggelaten wordt, kan alles zijn. En alle mogelijkheden zijn ook even waarschijnlijk.

Dat begrip *oneindig* blijft toch ook maar een raar ding. Er zijn oneindig veel mogelijkheden. Toch is er maar een relatief klein deel van echt mogelijk: alleen dat smalle stukje tussen $11,8$ V en $11,9$ V. Waarbij geldt dat $11,8$ V *wel* een mogelijke uitkomst is en $11,9$ *net niet*.

Hoe groot is eigenlijk de kans op *exact* $11,9$ V? Die is nul. Dat is ook weer zo vreemd met continue kansverdelingen. Doordat er oneindig veel waarden liggen in dat continue interval, is de kans op elke individuele uitkomst nul. Wiskundigen hebben het dan ook over een *kansdichtheid*. Er is namelijk wel een kans op een waarde tussen $11,80$ V en $11,81$ V. En er is ook een kans dat een waarde ligt tussen $11,80000$ V en $11,80001$ V. Terwijl daar ook weer oneindig veel getallen tussen liggen.

Je wordt een beetje duizelig van lang over getallen nadenken.

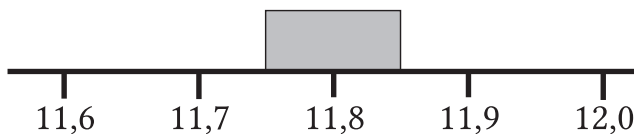


Die uniforme verdeling van de spanningen heeft een gemiddelde dat niet gelijk is aan 11,8 V, maar aan 11,85 V. Als je 11,8 V afleest, betekent dit dat de spanning gelijk is aan $11,85 \pm 0,05$ V.

Ook vreemd: om de gemeten waarde met zijn onzekerheid op juiste wijze te noteren, gebruik je één significant cijfer meer dan het display aangeeft. Er zijn geen drie, maar vier significante cijfers in de beste schatting. (En eentje in de onzekerheid.)

We zien dus dat het weglaten van alle volgende decimalen twee soorten fouten tot gevolg heeft. Een toevallige fout van 0,05 V én een systematische fout van 0,05 V. Tenzij je ervoor corrigeert. Als je weet dat de afronding altijd naar beneden is en je noteert de zojuist genoemde $11,85 \pm 0,05$ V, dan ben je dat probleem kwijt.

De kansdichtheidsfunctie ziet er na corrigeren (of voor een meetinstrument dat het goed afrondt) als volgt uit.



Hier zit geen systematische fout meer in. Er is een stuk dat we nooit zullen weten. Het enige wat we wel weten, is dat alle waarden even waarschijnlijk waren.

83. Systematische fout door toeval (2)

Een systematische fout kan ook uit een toevallige fout ontstaan als gevolg van de doorwerkingsformule. Bijvoorbeeld als je te maken hebt met een derde macht.

Voordat we dat gaan uitschrijven, even oefenen met haakjes uitwerken.

$$(a + b)^2 = (a + b) \times (a + b) = a^2 + 2ab + b^2$$

$$(a - b)^2 = (a - b) \times (a - b) = a^2 - 2ab + b^2$$

Deze kende je vast nog wel. Nu een macht hoger: geen kwadraat, maar tot de derde. We maken dan gebruik van de stap die we al hadden.

$$(a + b)^3 = (a + b)^2 \times (a + b) =$$

$$(a^2 + 2ab + b^2) \times (a + b) =$$

$$a^3 + 2a^2b + ab^2 + a^2b + 2ab^2 + b^3 = a^3 + 3a^2b + 3ab^2 + b^3$$

Nu gaan we deze uitgeschreven vorm gebruiken voor de oppervlakte van een cirkel en het volume van een bol, beide met een straal r .

Formules:

$$A = 4\pi r^2$$

$$V = \frac{4}{3}\pi r^3$$

Stel dat er in de straal een fout (*error*) zit ter grootte van Δr .

We nemen nu even een *fout* en niet een *onzekerheid*, en daarom ook een andere delta (de hoofdletter in plaats van de kleine letter). De reden daarvoor is dat het wat vreemd aanvoelt om te rekenen (vermenigvuldigen en delen) met een onzekerheid. Dat is een *bereik*, geen getal. Dat noteer je ook niet met een *plusteken*, maar met een *plusminteken*. Een fout is gewoon een waarde (die we meestal niet kennen). Je kunt de gevolgen van die (hypothetische, verzonnen) waarden doorrekenen.

Wat komt er uit ons sommetje als we niet met r rekenen, maar met $r + \Delta r$?

Dat kunnen we eenvoudig bekijken dankzij die formules die we net uitschreven, met voor a de waarde r en voor b de waarde Δr .

$$a^2 + 2ab + b^2$$

$$a^3 + 3a^2b + 3ab^2 + b^3$$

Ingevuld en vermenigvuldigd met de factor die ervoor staat:

$$A = 4\pi(r^2 + 2r\Delta r + (\Delta r)^2)$$

$$V = \frac{4}{3}\pi(r^3 + 3r^2\Delta r + 3r(\Delta r)^2 + (\Delta r)^3)$$

De fout Δr kan positief of negatief zijn. We gaan uit van een symmetrisch verdeelde fout (normaal of uniform, meestal). Die symmetrie raken we kwijt door het doorrekenen. Vul in plaats van Δr maar eens $-\Delta r$ in.

$$A = 4\pi(r^2 - 2r\Delta r + (\Delta r)^2)$$

$$V = \frac{4}{3}\pi(r^3 - 3r^2\Delta r + 3r(\Delta r)^2 - (\Delta r)^3)$$

Als de fout in de oppervlakte van een cirkel en het volume van een bol symmetrisch zou doorwerken, dan zou de oppervlakte van een cirkel met straal $r + \Delta r$ (in positieve zin) net zoveel moeten afwijken van de oppervlakte van een cirkel met straal $r - \Delta r$ (in negatieve zin). Een fout “de ene kant op” moet net zoveel invloed hebben als een fout “de andere kant op”. Alleen dan heffen positieve en negatieve fouten elkaar op.

We zouden het verder kunnen uitschrijven door de oppervlakte en het volume zonder de fout Δr ervan af te trekken. Dan krijg je dit.

$$\Delta A = 4\pi(-2r\Delta r + (\Delta r)^2)$$

$$\Delta V = \frac{4}{3}\pi(-3r^2\Delta r + 3r(\Delta r)^2 - (\Delta r)^3)$$

Om het nog leuker te maken, zou je die verschillen kunnen berekenen voor een positieve én een negatieve fout van dezelfde grootte. Daar zou in het ideale geval nul uitkomen. Maar dat is niet zo.

Dat zul je altijd zien: gevallen zijn maar zelden ideaal. Ja, in films. En in boeken. En in de verhalen van docenten als ze voor de klas staan. En tijdens diploma-uitreikingen, want dan is opeens alles goed. Maar nu dus niet. Meestal niet trouwens, let er maar eens op.

We nemen nu voor altijd een positieve waarde en vullen die in de formule in met een plus of een min ervoor.

$$\Delta A_{positief} = 4\pi(-2r\Delta r + (\Delta r)^2)$$

$$\Delta A_{negatief} = 4\pi(2r\Delta r + (\Delta r)^2)$$

$$\Delta V_{positief} = \frac{4}{3}\pi(-3r^2\Delta r + 3r(\Delta r)^2 - (\Delta r)^3)$$

$$\Delta V_{negatief} = \frac{4}{3}\pi(3r^2\Delta r + 3r(\Delta r)^2 + (\Delta r)^3)$$

Nu kunnen we ze van elkaar aftrekken.

$$\Delta A_{negatief} - \Delta A_{positief} =$$

$$4\pi((2r\Delta r + (\Delta r)^2) - (-2r\Delta r + (\Delta r)^2)) =$$

$$16\pi r\Delta r$$

$$\Delta V_{negatief} - \Delta V_{positief} =$$

$$\frac{4}{3}\pi((3r^2\Delta r + 3r(\Delta r)^2 + (\Delta r)^3) - (-3r^2\Delta r + 3r(\Delta r)^2 - (\Delta r)^3)) =$$

$$\frac{4}{3}\pi(6r^2\Delta r + 2(\Delta r)^3) =$$

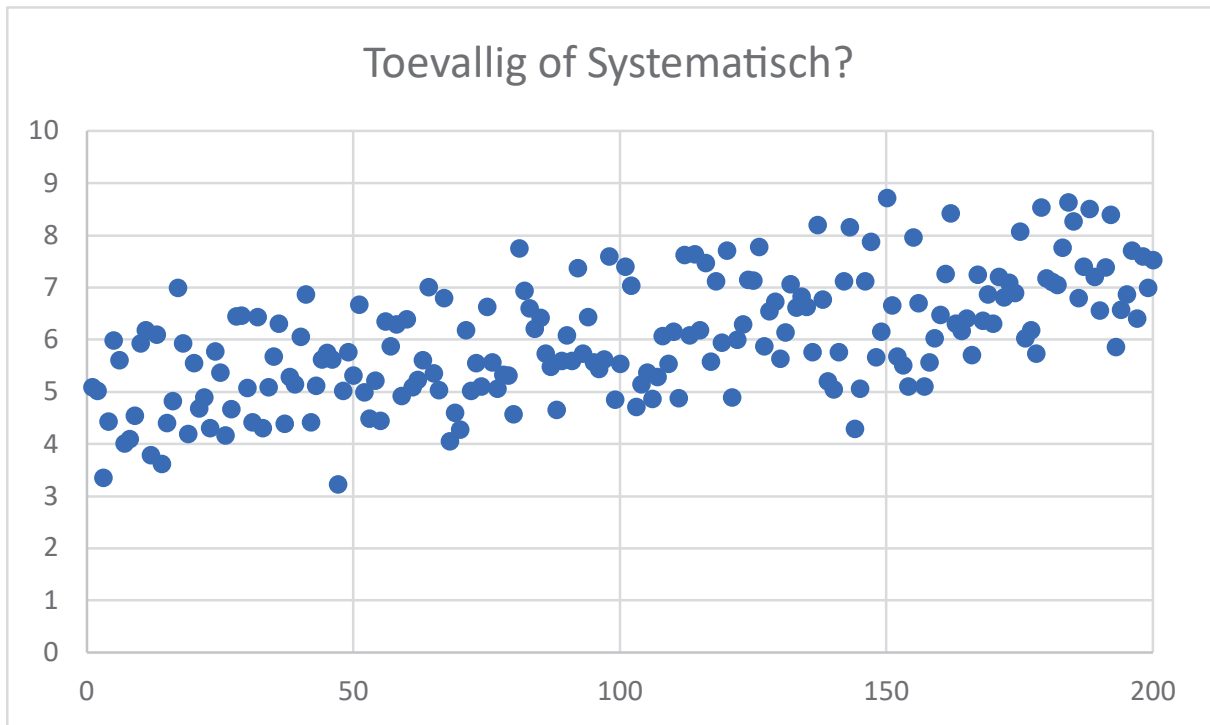
$$\frac{8}{3}\pi(3r^2\Delta r + (\Delta r)^3) =$$

We zien dus dat er zowel bij de oppervlakte van de cirkel als bij het volume van de bol een altijd positieve term overblijft: r en Δr zijn immers positief. De positieve term is groter naarmate de fout groter is. Voor beide geldt trouwens ook dat hij groter wordt naarmate de straal zelf groter is. Best opmerkelijk... De doorwerking van de *absolute fout* wordt dus groter als de *grootheid* zelf ook groter is. Het klopt, maar het is niet vanzelfsprekend.

Berendsen (2011) gebruikt in bijlage A2 van zijn boek ook het voorbeeld van het volume van een bol. Niks mis mee, maar hij had ook de oppervlakte van een cirkel kunnen nemen. Het zit namelijk niet in die derde macht.

84. Systematisch toenemende fout

Onderstaande grafiek toont de uitkomst van tweehonderd metingen aan een denkbeeldige grootheid.



Er staan geen eenheden bij. Het zou een spanning kunnen zijn in volt of millivolt, een temperatuur in graden Celsius, of wat dan ook. Voor het gemak kijken we nu alleen even naar de getallen.

De metingen zijn gedaan in een periode van één uur. Precies om de achttien seconden werd er (automatisch) een meting verricht en omdat $200 \times 18 \text{ s} = 3600 \text{ s}$ is er dus exact een uur gemeten.

De vraag is: kunnen we iets zeggen over de toevallige en/of systematische fout in de metingen?

Het moge duidelijk zijn dat er een flinke spreiding in de metingen zat. De laagste waarde was 3,23, de hoogste 8,71. Het gemiddelde was 6,02 en de standaardafwijking was 1,14. Je zou dus kunnen zeggen dat de uitkomst van de meting gelijk was aan $6,0 \pm 1,1$.

Bij het kijken naar de grafiek valt er iets op. De spreiding in de waarden is vrij groot, maar er lijkt toch een bepaalde trend zichtbaar te zijn. Gemiddeld genomen nemen de waarden toe in de loop van dat ene uur waarin er gemeten is. De waarden liggen rechts in de grafiek gemiddeld hoger dan links.

Excel kan ons een handje helpen met een trendlijn, maar die hebben we hier weggelaten. Als je die zou trekken, zag je inderdaad een stijgende lijn.

Omdat we de fictieve metingen zelf hebben gegenereerd, weten we hoe ze tot stand zijn gekomen. De tweehonderd getallen bestaan uit normaal verdeelde ruis met een gemiddelde van 0 en een standaardafwijking van 1, opgeteld bij een lineaire functie.

$$y(t) = 5 + 0,01t$$

Als het een temperatuurmeting was, zou je dus kunnen zeggen dat de temperatuur lineair toenam van 5,0 °C naar 7,0 °C. Een toename die gelijk is aan tweemaal de standaardafwijking van de ruis, de toevallige fout die erbovenop zat.

Kun je nu concluderen dat er sprake is van een systematische fout die van 0,0 °C naar 2,0 °C oploopt? Nee. Het is in ieder geval niet de enige mogelijkheid. Misschien is er wel een systematische fout die van 1,3 °C naar 3,3 °C oploopt. Of van -0,5 °C naar 1,5 °C.

De *verandering* in de *systematische* fout is gedeeltelijk te vinden door *statistische* analyses. En daarmee geldt dus hetzelfde als voor toevallige fouten. Veranderingen kun je waarnemen en analyseren, vaste afwijkingen zijn ‘ingekapseld’ in de metingen zelf.

Chauvenet

85. Het criterium van Chauvenet

Soms heb je een meting die nogal ver afwijkt van de rest. Dat voedt het vermoeden dat die meting niet goed is, maar hoe weet je dat?

William Chauvenet (24 mei 1820 – 13 december 1870) was een professor in de wiskunde, astronomie, navigatie en landmeetkunde. Hij bedacht een criterium op grond waarvan je – op basis van statistiek – kon besluiten of je een waarde zou mogen verwerpen.

Simpel gezegd komt zijn methode hierop neer.

1. Schat het gemiddelde en de standaardafwijking van de normale verdeling die ten grondslag ligt aan de metingen
2. Bepaal de kans dat de ‘verdachte waarde’ optreedt.
3. Bereken de verwachtingswaarde van hoe vaak zoiets gebeurt bij het aantal metingen dat je hebt.
4. Verwerp de meting als die verwachtingswaarde kleiner is dan een half. Een half zit immers dichterbij nul dan bij één.

Hoe werkt het criterium van Chauvenet?

Je veronderstelt dat de metingen (inclusief de verdachte waarde!) afkomstig zijn van een normale verdeling. Wat je vervolgens wilt weten, is hoe die ene uitschieter (of een grotere) voorkomt in dit aantal metingen. Beter gezegd: wat de verwachtingswaarde is van het aantal keren dat die waarde voorkomt. Als dat minder is dan een half (0,5), dan verwerp je de meting. De verwachtingswaarde zit dan namelijk dichterbij 0 dan bij 1.

Het is belangrijk dat je bij de berekening uitgaat van de juiste formule voor de standaardafwijking, namelijk die voor een *steekproef*, met de $N - 1$ in de noemer.

Taylor gebruikt in zijn boek het voorbeeld van iemand die slingertijden meet. De gevonden periodetijden (in seconden) zijn als volgt:

3,8, 3,5, 3,9, 3,9, 3,4, 1,8.

In deze reeks zijn de getallen eenvoudig. Afgerond op één decimaal vind je:

$$\bar{x} = 3,4 \text{ s}$$

$$s_x = 0,8 \text{ s}$$

De verdachte waarde van 1,8 s wijkt precies twee keer de geschatte standaardafwijking af van het geschatte gemiddelde. Vermoedelijk heeft Taylor die waarde opzettelijk zo gekozen, want het verschil tussen 3,4 en 1,8 is precies twee keer die 0,8. En de kans dat een waarde minstens twee keer de standaardafwijking afwijkt van het gemiddelde is bij een normale verdeling ongeveer 5%. Je kunt dus uit je hoofd uitrekenen dat de verwachtingswaarde van het aantal metingen met een dergelijke afwijking gelijk is aan $0,05 \times 6 = 0,3$. En dat is minder dan 1.

Het is een beetje vreemd om statistiek te bedrijven met zulke kleine aantallen. Zes metingen is niet echt veel. En bij de bepaling van het gemiddelde en de standaardafwijking neem je ook nog eens de verdachte waarde mee! In het toetsen van hypothesen is dat gebruikelijk: je bewijst dat iets onwaarschijnlijk is door eerst aan te nemen dat het waar is en daarna te bepalen dat de kans dit die aanname klopt klein is. Als je de verdachte waarde niet mee zou nemen, werk je eigenlijk al naar het antwoord toe. Je stopt je eigen vooroordeel in het verwerpen.

Wat nu als er meer metingen geweest waren? Je verworpt een verdachte meting als de verwachtingswaarde kleiner is dan 0,5, maar als je meer metingen hebt, wordt de verwachtingswaarde toch ook groter?

We gaan nu, in plaats van hierover na te denken, een sommetje maken. We laten de vijf ‘gewone’ metingen allemaal twee keer voorkomen en die ene verdachte meting één keer, door die eerste vijf waarden achter de reeks te plakken. De metingenserie is dan (met weglating van de eenheden):

3,8, 3,5, 3,9, 3,9, 3,4, 1,8, 3,8, 3,5, 3,9, 3,9, 3,4.

Dat resulteert in een ander gemiddelde en een andere standaardafwijking: 3,5 en 0,6. De kans op een afwijking van 1,9 of hoger is lastiger uit te rekenen, omdat het niet meer exact een factor 2 maal de standaardafwijking is: $1,6/0,8$ was precies 2, maar $1,7/0,6 \approx 2,83$. De overschrijdingskans is al een stuk kleiner, namelijk 0,73 %. De verwachtingswaarde is 0,08: veel minder dan 0,5.

Als je niet weet hoe je een overschrijdingskans bij een normaal verdeeld proces berekent: in een volgend hoofdstuk wordt dat uitgelegd.

We gaan nu de tussenstappen van de berekening in een tabel zetten, met lekker veel cijfers achter de komma. Dat doen we voor de oorspronkelijke zes metingen (vijf plus één verdachte), maar ook voor een set van elf metingen (twee keer die vijf plus één verdachte) en 21 metingen (vier keer die vijf plus

één verdachte). We voegen bovendien een extra kolom toe, namelijk eentje waar 101 metingen gedaan zijn: twintig keer het setje van die vijf ‘goede’ metingen en die ene verdachte ertussendoor.

Dit doen we dit? Om wat meer inzicht en gevoel te krijgen van de invloed die het aantal waarnemingen heeft op het al dan niet verwerpen van de uitschieter. Het klinkt aannemelijk dat je bij een groot aantal waarnemingen die ene verdachte waarde eerder durft te verwerpen dan bij een klein aantal, omdat die verdachte waarde meedoet in de bepaling van het gemiddelde en de standaardafwijking. Die worden bij kleine aantal sterker ‘gekleurd’ door de uitschieter. Bij honderd metingen vormt die ene vreemde eend in de bijt maar één procent van het geheel. Dat zie je nauwelijks meer terug in $\bar{x}$ en s_x .

In de tabel staat bij alle variabelen ook hun betekenis weergegeven. Het voorbeeld helpt dus ook als je beter wilt snappen hoe het werkt.

$\bar{x}$	gemiddelde van de steekproef, inclusief verdachte waarde
s_x	standaardafwijking van de steekproef (inclusief...)
x_v	verdachte waarde
$ \bar{x} - x_v $	¹ afwijking (verschil gemiddelde en verdachte waarde)
t_v	¹ afwijking gedeeld door de standaardafwijking
P_v	kans op een overschrijding in de normale verdeling
N	aantal metingen
n	² verwachtingswaarde van aantal keer zo’n uitschieter

¹ Normering maakt er een standaardnormale verdeling van.

² Zo’n uitschieter: net zo groot of groter, beide kanten op.

	6 metingen	11 metingen	21 metingen	101 metingen
$\bar{x}$	3,3833	3,5273	3,6095	3,6812
s_x	0,8035	0,6101	0,4647	0,2824
x_v	1,8	1,8	1,8	1,8
$ \bar{x} - x_v $	1,5833	1,7273	1,8095	1,8812
t_v	1,9705	2,8313	3,8943	6,6617
P_v	0,0488	0,0046	$9,84 \cdot 10^{-05}$	$2,71 \cdot 10^{-11}$
N	6	6	6	6
n	0,2927	0,051	0,0021	$2,73 \cdot 10^{-9}$

De eenheden hebben we weggelaten. De bovenste vier zijn in seconden, de onderste vier zijn dimensieloos.

In de laatste tabel wordt zichtbaar wat het belang van veel metingen is als je statistiek bedrijft. De verwachtingswaarde op grond waarvan er besloten wordt, neemt af van 0,3 naar 0,05 als je van zes naar elf metingen gaat, en naar 0,002 bij 21 metingen. De kans op zo'n afwijkende waarde bij 101 metingen is extreem klein. Die 1,8 kan in principe nog steeds géén uitschieter zijn, maar die kans is drie op miljard. Dat mag je verwaarlozen.

86. Chauvenet in andere boeken

Voordat Taylor (1997) het in zijn boek over het verwerpen van metingen gaat hebben, staat hij erbij stil dat deze nare vraag (*awkward question*) in wetenschappelijke kringen *controversieel* is: aanvechtbaar, omstreden.

This chapter discusses the awkward question of whether to discard a measurement that seems so unreasonable that it looks like a mistake. This topic is controversial; some scientists would argue that discarding a measurement just because it looks unreasonable is never justified.

Waarom zou je metingen verwerpen? Je hebt ze toch gedaan? De andere kant is ook vreemd: als je het sterke vermoeden hebt dat iets niet klopt, dan kun het toch moeilijk voor waar houden?

Er is nog een aspect waar vaak aan voorbij gegaan wordt. Als er iets fout gaat bij het meten, *kan* dat resulteren in een uitschieter. Die is dan verdacht en kan om die reden nader worden onderzocht. Maar een meetfout kan natuurlijk ook onopvallend tussen de rest verscholen blijven, doordat het ‘toevallig’ geen uitschieter is. Dan geeft het ook niet, zou je zeggen. Maar stel nu eens dat de helft van de metingen fout is en telkens bijna dezelfde waarde oplevert. Als die waarde dicht bij het gemiddelde ligt, zal hij daar weinig invloed op hebben. Maar op de standaardafwijking heeft hij die wél! Zo zouden verkeerde metingen een te kleine onzekerheid kunnen suggereren.

Het voorbeeld dat Taylor gebruikt, is eigenlijk raar. Als we ervan uitgaan dat de eerste vijf metingen ‘goed’ zijn, liggen alle waarden tussen 3,4 s en 3,9 s. Dat is dan een halve seconde verschil tussen de langste en de kortste slinger-tijd. Een spreiding die waarschijnlijk veroorzaakt wordt door het menselijke onvermogen om precies op het goede moment de stopwatch in te drukken. Waarom heb je dan geen *tien* herhalingen gedaan? Dan wordt die onzekerheid in het aflezen gedeeld door tien, en valt het allemaal wel mee.

Even tussendoor: wat nu als die 1,8 s de enige goede meting was, en alle andere waren verkeerd?

Als de eerste vijf metingen correct waren, is de gemiddelde slingertijd 3,7 s. Je meetresultaat zou dan $3,7 \pm 0,2$ s zijn. Die 1,8 s die ook gemeten werd, wijkt ruim de helft af van dat gemiddelde. Je voelt wel aan dat dit niet kan kloppen, vooral ook omdat je weet wat je aan het meten bent.

Je kent de natuurkundige formule: $T_0 = 2\pi \sqrt{\frac{l}{g}}$

Het lijkt me sterk dat g veranderd is, laat staan π ... En om die term in de wortel te halveren, moet wat onder het wortelteken staat een factor 4 kleiner worden. Dan zou die slinger tot een kwart van z'n lengte gekrompen zijn.

Bevington & Robinson (2002) vertellen in één zin van drie regels hoe het zit.

The established condition for discarding data in such circumstances is known as Chauvenet's criterion, which states that we should discard a data point if we expect less than half an event to be farther from the mean than the suspect point.

Vrij vertaald in het Nederlands:

Verwerp een waarde als het verwachte aantal metingen dat verder van het gemiddelde verwijderd ligt dan die waarde, kleiner is dan een half.

Eigenlijk is hiermee alles gezegd. Minder dan een halve pagina, en dan nog tweetalig ook. Maar mocht je hier niet voldoende aan hebben, ga dan naar het volgende hoofdstuk. Als je wilt weten wat Bevington & Robinson nog meer te melden hebben over de verdwijntruc van Chauvenet, lees dan de rest van deze bladzijde ook.

Removing an outlying point has a greater effect on the standard deviation than on the mean of a data sample, because the standard deviation depends on the squares of the deviations from the mean. Deleting one such point will lead to a smaller standard deviation and perhaps another point or two will now become candidates for rejection. We should be very cautious about changing data unless we are confident that we understand the source of the problem we are seeking to correct, and repeated point deletion is generally not recommended. The importance of keeping good records of any changes to the data sample must also be emphasized.

Vertaald: Verwijder niet zomaar een meetwaarde en leg goed vast wat je doet.

(Als je veel weglaat, is zo'n Nederlandse vertaling toch korter.)

Oefening

Schrijf een Pythonprogramma waarmee je via een simulatie bepaalt hoe groot de kans is dat je in een normaal verdeelde set van tien metingen ten onrechte een uitschieter verwerpt.

87. Overschrijdingskansen

Je weet vast wel wat een gemiddelde en een standaardafwijking is, en snapt ongetwijfeld waarom je niet de μ en de σ , maar de $\bar{x}$ en de s gebruikt. Dan heb ik goed nieuws voor je: de eerste oefening onderaan dit hoofdstuk wordt een makkie voor je.

Ook een verdachte meting is niet zo moeilijk: dat er eentje wat afwijkt, dat zie je zo. Maar dat die ene waarde zoveel afwijkt dat hij verworpen moet worden, daar zul je eerst nog aan moeten rekenen.

Het gaat om de kans dat minimaal die waarde voorkomt, als je veronderstelt dat er sprake is van een normaal verdeeld proces met een gemiddelde van μ en een standaardafwijking σ .

Ai... Die zouden we niet gebruiken toch?

Nee. We zouden het best willen, maar we *kennen* de μ en de σ niet. (Als je het antwoord op Oefening 1 niet wist, dan weet je het nu...) Daarom gebruiken we *bij gebrek aan beter* de $\bar{x}$ en de s : de standaardafwijking en het gemiddelde uit de steekproef. Je moet toch wát...

Hoe groot is de kans dat een getal z de uitkomst is van een normaal verdeeld proces met een bepaalde μ en de σ . Het antwoord zal lezers die weinig verstand hebben van statistiek verbazen. Die kans is namelijk nul...

Nul? Ja. De kans op precies z is 0. Een normaal verdeeld proces heeft oneindig veel mogelijk uitkomsten. De kans op elk van die uitkomsten is 0. Wat we wel kunnen doen, is de kans berekenen dat een waarde wordt gevonden in een bepaald bereik. Of dat er een waarde wordt gevonden die groter of kleiner is dan z .

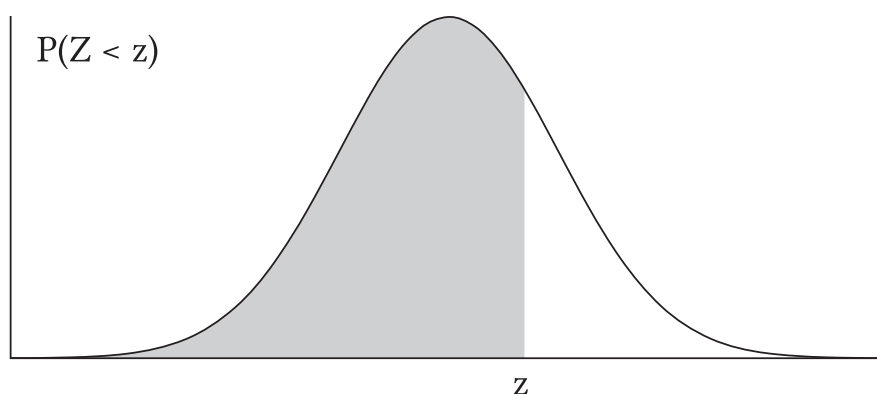
Eigenlijk klopt de formulering niet. Het gaat niet om de kans dat die ene uitschieter voorkomt, maar dat een uitschieter minstens zo veel afwijkt van het gemiddelde. Die kun je bepalen door de overschrijdingskansen te berekenen.

Ter illustratie kijken we naar een normaal verdeelde grafiek, met daarin een waarde z . De totale oppervlakte onder de kromme is gelijk aan 1. Dat is namelijk de som van alle kansen op alle mogelijke waarden. De kans op “iets” is nu eenmaal 100%. De oppervlakte van het grijze gebied is de kans dat er een waarde optreedt die kleiner is dan z . De oppervlakte van het witte gebied is de kans dat er een waarde optreedt die groter is dan z .

Hé, maken we hier geen fout? Als ik wil weten hoeveel studenten op de opleiding jonger zijn dan 20 jaar, kan ik ze in twee groepen verdelen. De studenten met een leeftijd kleiner dan 20 jaar en de studenten met een leeftijd groter dan 20 jaar. Maar dan mis ik een flinke groep, namelijk de twintigjarigen! Die zullen meer dan 10% van de populatie vormen. We moeten de grenzen dus anders trekken. Er zijn studenten die *jonger dan 20 jaar* zijn en studenten die *20 jaar of ouder* zijn. De groepen < 20 en ≥ 20 vormen samen iedereen.

Waarom doen we daar dan bij die z niet moeilijk over? Omdat de kans op exact z toch gelijk is aan 0. De jongste student zal 16 of 17 zijn en de oudste ergens rond de 30. Het aantal verschillende leeftijden is heus geen twintig. In ieder geval veel minder dan dat oneindige aantal mogelijkheden uit ons sommetje.

De kans dat een willekeurige (toevallige, stochastische) Z kleiner is dan z is gelijk aan de oppervlakte van het grijze gebied in Figuur 87.1.



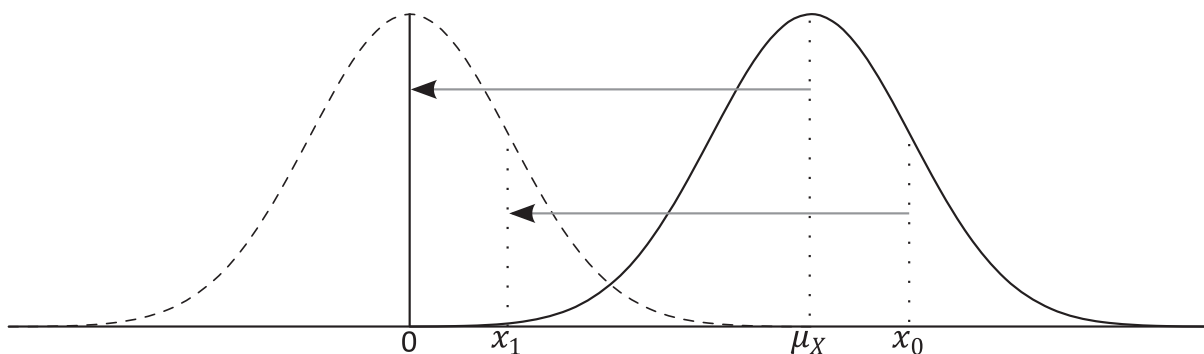
Figuur 87.1 De kans op een waarde z of kleiner

Hoe groot is die oppervlakte? Meer dan 0,5 zo te zien. Maar hoeveel precies?

De klassieke manier om dat uit te vinden is: aflezen uit een tabel die al die kansen weergeeft voor een standaardnormale verdeling.

Maar wacht even... De waarde z komt toch niet uit een standaardnormale verdeling? Een standaardnormale verdeling is een gewone normale verdeling, maar dan een heel bijzondere. Die heeft namelijk een μ en een σ die lekker makkelijk zijn: $\mu = 0$ en $\sigma = 1$. Met twee slimme trucjes kunnen we elke normale verdeling omtoveren in een standaardnormale verdeling. Je trekt eerst het gemiddelde μ_x van jouw “eigen” verdeling af van alle waarden, en je deelt die vervolgens door “jouw” σ_x .

Die twee stappen kun je in een plaatje weergeven. Door van alle waarden μ_x af te trekken, verschuift het gemiddelde naar 0 en krijgen we een normale verdeling die symmetrisch is rondom 0.

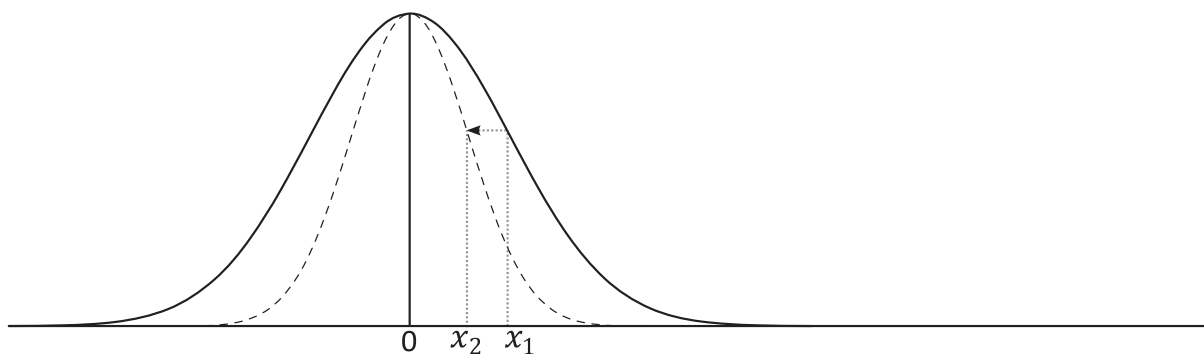


Figuur 87.2 "opschuiven" van de verdeling zodat het gemiddelde 0 wordt

Wat gebeurt er dan met een willekeurige waarde x_0 ? Die wordt met μ_X verlaagd. In formulevorm:

$$x_1 = x_0 - \mu_X$$

Door alle waarden door σ_X te delen, wordt de grafiek (in dit geval) smaller en wordt de standaardafwijking gelijk aan 1.



Figuur 87.3 "samendrukken" van de verdeling om een standaardafwijking van 1 te krijgen

Wat gebeurt er dan met onze waarde x_1 ? Ook die wordt gedeeld door σ_X . In formulevorm:

$$x_2 = \frac{x_1}{\sigma_X} = \frac{x_0 - \mu_X}{\sigma_X}$$

Wiskunde is ontzettend leuk! Maar af en toe word je horendol van al die tekenjjes en indexjes en streepjes en bocheltjes. Waarom staat hier nou soms een hoofdletter X en soms een kleine letter x ? Omdat die twee niet hetzelfde zijn. De hoofdletter X duidt een *kansvariabele* aan. (Een stochastische variabele noemen ze dat ook wel.) Die kan verschillende waarden aannemen, volgens een bepaalde kansverdeling. De kleine letter x is één bepaalde waarde. Dat is dus geen kansvariabele meer. Hij is gewoon 1,2 of 1,8 of $-0,4$ of welk ander getal dan ook. Die indexjes 0, 1 en 2 zijn er alleen maar om de drie waarden uit elkaar te houden.

Dankzij dat opschuiven en delen is de kansdichtheidsfunctie genormeerd tot een *standaardnormale verdeling*. En voor die bijzondere normale verdeling met $\mu = 0$ en $\sigma = 1$ vind je tabellen in wiskundeboeken. En in dit boek.

Nu terug naar het voorbeeld. Hoe groot was onze verdachte waarde? Hoe groot was ons gemiddelde en hoe groot was de standaardafwijking?

x_{ver}	1,8
$\bar{x}$	3,3833
s_x	0,8035

Zoals gezegd: we kennen het echte gemiddelde en de echte standaardafwijking niet. Daarom gebruiken we de schattingen uit de steekproef.

$$x_{norm} = \left| \frac{x_{ver} - \bar{x}}{s_x} \right| \approx \left| \frac{1,8 - 3,3833}{0,8035} \right| \approx 1,9705$$

Deze genormeerde x_{norm} is precies wat we bij Chauvenet aanduiden als “het aantal keren de standaardafwijking dat de verdachte waarde afwijkt van het gemiddelde”. En we willen weten hoe groot de kans is dat zo’n afwijking zo veel keer (of meer) voorkomt.

Hoe lees je dat af uit die tabel? We hoeven niet alle rijen te bekijken, alleen de waarden die rond de 2 zitten.

In de tabel vind je de kansen in vier decimalen (vier significante cijfers) voor getallen met twee decimalen (maximaal drie significante cijfers). Daarom moeten we 1,9705 afronden op 1,97. Die 1,9 vind je in de linker kolom. De tweede decimaal, die 7 dus, vind je in de bovenste rij. We zijn op zoek naar het getal dat in de rij staat die met 1,9 begint en de kolom waar 7 boven staat. Dat getal is 0,9756.

Wat is de betekenis daarvan ook alweer? Het is de kans dat één waarde van een standaardnormaal verdeeld proces kleiner is dan 1,97. Die kans is dus 97,56%. De kans dat die waarde groter is, is dus 2,44%. Maar let op: we zijn geïnteresseerd in uitschieters *naar beide kanten*. We moeten dus ook de gevallen meenemen waar die ene waarde kleiner is dan $-1,97$. De gevonden kans moet daarom nog met 2 worden vermenigvuldigd. Maar dan zijn we er!

$$P(X > 1,97) = 2 \times (1 - 0,9756) = 2 \times 0,0244 = 0,0488$$

88. Overschrijdingskansen online

Naast de klassieke aanpak met het omrekenen naar standaardnormaal en het aflezen in een tabel en de werkbladbenadering van *Excel* in het volgende hoofdstuk, is er ook nog het gebruik van *online* hulpmiddelen. Zoek met Google op ‘online normale verdeling’ en je vindt webpagina’s waar je drie getallen invult en de oplossing hebt. Met alleen het gemiddelde, de standaardafwijking en de verdachte waarde ben je er al.

Bijna dan. Chauvenet is gebaseerd op verwachtingswaarden. En de overschrijdingskans gaat uit van het dubbele, omdat de kans op een uitschieter aan een van beide kanten bepaald wordt.

Normal Distribution Calculator

find the area under normal distribution curve

show help ↓↓ examples ↓↓

If X is a normally distributed variable with mean $\mu =$ and standard deviation $\sigma =$ find one of the following probabilities:

☐ $P(\text{ } < X < \text{ })$

☐ $P(X > \text{ })$

☒ $P(X < \text{ } 1.8 \text{ })$

☐ Hide steps

3.38333333
0.803533862

Area = ?

1.8 $\mu = 3.38333333$

Figuur 88.1 Voorbeeld van een online hulpmiddel om overschrijdingskansen te berekenen

Let bij de *online* tools goed op of je een punt of een komma gebruikt.

Oefening

Zoek *online* twee verschillende hulpmiddelen om overschrijdingskansen te berekenen en controleer daarmee de Chauvenetberekening uit voorgaande hoofdstukken.

89. Chauvenet in Excel

Moet het nou echt zo moeilijk? Nee. Vroeger moest het zo, toen we nog geen computers hadden. Tegenwoordig doe je het hooguit om meer inzicht in de lesstof te krijgen. Hoe pak je dit aan in *Excel*?

`=2*NORM.VERD(1,8;3,8333;0,7661;WAAR)`

En klaar ben je! Volgens *Excel* is de uitkomst 0,04878. In vier significante cijfers net niet hetzelfde als wij “met de hand” gevonden hadden. In drie significante cijfers is het exact hetzelfde.

In plaats van het invullen van die harde getallen, kun je in *Excel* ook met velden werken. Dan komt de formule er zo uit te zien.

`=2*NORM.VERD(C8;C1;C2;WAAR)`

Die F4, F1 en E7 zijn de velden in mijn werkblad waar de formules staan voor de verdachte waarde, het gemiddelde en de standaardafwijking. WAAR betekent dat we cumulatief werken. Nu vinden we 0,048785389886591.

Is het nou niet een beetje overdreven om *velden* op te nemen in de formule in plaats van *getallen*? En een antwoord te geven met alle decimalen waarmee *Excel* rekt? Om met dat laatste te beginnen: ja, dat is overdreven. Maar het rekenen met velden in plaats van getallen verdient wel echt de voorkeur. Niet alleen omdat er geen reden is om decimalen weg te laten, maar vooral ook omdat het werken met formules ervoor zorgt dat je in de oorspronkelijke dataset aanpassingen kunt maken die meteen doorwerken.

In onderstaande tabel staat de inhoud van een *Excel*-werkblad met dertien rijen (1 t/m 13) en vier kolommen (A t/m D). De kolommen bevatten:

- A. De gegevens waarop de analyse wordt uitgevoerd.
- B. Een tekst om aan de gebruiker te laten zien wat er op die rij wordt berekend.
- C. De uitkomst van die berekening. Er worden nul, een of twee decimalen getoond, afhankelijk van de waarde die er staat.
- D. De uitkomst van die berekening, maar dan in alle decimalen die *Excel* standaard toont.

In kolom D staan nog steeds niet alle decimalen die *Excel* heeft. De uitkomst van de verwachtingswaarde is 0,292712339319546. *Excel* rekt namelijk met vijftien decimalen.

In het kader rechts staan de gebruikte formules. Het gemiddelde (1), de standaardafwijking van de steekproef (2), de maximale waarde (3) en de minimale

waarde (4) worden allemaal gehaald uit kolom A. Dat er tussen de haken A:A staat in plaats van A1:A6 zorgt ervoor dat je gegevens kunt toevoegen of verwijderen zonder dat je deze formules hoeft aan te passen. Let op dat je STDEV.S neemt en niet STDEV.P. (Geen *populatie* maar *steekproef*.)

Cellen C5, C6 en C7 zijn eenvoudige berekeningen op C3 en C4. C8 is een bijzondere functie die een lijkt op de vraagtekenoperator in de taal C. Je test of C5 groter is dan C6. Als dat waar is, is de uitkomst gelijk aan C5. Als dat onwaar is, is het gelijk aan C6. Zo vind je de grootste waarde van de twee.

Met AANTALARG wordt geteld hoeveel metingen er in kolom A staan.

	A	B	C	D	
1	3,8	gemiddelde	3,38	3,383333333	=GEMIDDELDE(A:A)
2	3,5	standaardafwijking	0,80	0,803533862	=STDEV.S(A:A)
3	3,9	hoogste	3,9	3,9	=MAX(A:A)
4	3,9	laagste	1,8	1,8	=MIN(A:A)
5	3,4	afwijking 1	0,52	0,516666667	=C3-C1
6	1,8	afwijking 2	1,58	1,583333333	=C1-C4
7		grootste afwijking	1,58	1,583333333	=MAX(C5:C6)
8		verdachte waarde	1,8	1,8	=ALS(C5>C6;C3;C4)
9		aantal metingen	6	6	=AANTALARG(A:A)
10		aantal keer standaardafwijking	1,97	1,97046249	=C7/C2
11		kans kleiner / groter	0,0241	0,024392695	=NORM.VERD(C8;C1;C2;WAAR)
12		kans grotere afwijking	4,8%	0,04878539	=2*C11
13		verwachtingswaarde	0,29	0,292712339	=C9*C12

Oefeningen

Maak een werkblad in *Excel* waar in kolom A de zes gegevens staan uit het voorbeeld van dit hoofdstuk: 3,8 3,5 3,9 3,9 3,4 1,8.

1. Vul kolom B met formules die nodig zijn voor de *rough and ready approach* van Hughes & Hase.
2. Voeg aan kolom B velden toe waarmee je het maximale verschil met de mediaan berekent.

90. Chauvenet in Python

Dezelfde analyse die we in *Excel* hebben gedaan, kan natuurlijk ook met een handige programmeertaal als *Python*.

Het voordeel van *Excel* is dat het sneller gaat; programmeren in een taal is meestal wat meer werk dan het invullen van rijen en kolommen.

Het voordeel van *Python* is dat het krachtiger en flexibeler is. Je kunt een functie maken en die bijvoorbeeld in een lus aanroepen.

De code van een Pythonprogramma dat de berekeningen doet, staat op de volgende pagina. Daarbij zijn tussenstappen in twee of vier decimalen afgerond. Normaal gesproken is dat niet gewenst, maar we willen de berekening kunnen vergelijken met een handmatige aanpak, waarin gebruik wordt gemaakt van een tabel van de standaardnormale verdeling. Om exact daarmee in overeenstemming te zijn, moet je dezelfde significantie hanteren.

Python blijkt met één of twee decimalen meer te rekenen. Over die laatste cijfers zijn *Python* en *Excel* het onderling niet altijd eens.

Tabel 90.1 Onafgeronde waarden in *Excel* en *Python* voor Chauvenetberekening

berekende waarde	Excel	Python
gemiddelde	3,38333333333333	3,38333333333333 <u>3</u>
standaardafwijking ("N-1")	0,80353386155573 <u>1</u>	0,80353386155573 <u>2</u> <u>3</u>
hoogste	3,9	3,9
laagste	1,8	1,8
afwijking 1	0,51666666666666 <u>7</u>	0,51666666666666 <u>6</u>
afwijking 2	1,58333333333333	1,58333333333333 <u>3</u>
grootste afwijking	1,58333333333333	1,58333333333333 <u>3</u>
verdachte waarde	1,8	1,8
aantal metingen	6	6
aantal keer standaardafwijking	1,9704624896178 <u>3</u>	1,9704624896178 <u>2</u> <u>5</u> <u>6</u>
kans kleiner (of groter)	0,02439269494329 <u>5</u>	0,02439269494329 <u>5</u> <u>6</u> <u>5</u> <u>6</u>
kans grotere afwijking	0,04878538988659 <u>1</u>	0,04878538988659 <u>1</u> <u>3</u> <u>1</u>
verwachtingswaarde	0,292712339319547	0,292712339319547 <u>9</u>

```

import numpy as np
import statistics as st
from scipy.stats import norm

# rij met meetwaarden
waarden = [3.8, 3.5, 3.9, 3.9, 3.4, 1.8]

# statistische berekeningen
gemiddelde = round(np.mean(waarden), 2)
standaardafwijking = round(st.stdev(waarden), 2)
hoogste = np.amax(waarden)
laagste = np.amin(waarden)
afwijking_hoogste = round(hoogste - gemiddelde, 2)
afwijking_laagste = round(gemiddelde - laagste, 2)
afwijking_max = max(afwijking_hoogste, afwijking_laagste)
aantal_waarden = len(waarden)

# verdachte waarde
if afwijking_hoogste > afwijking_laagste:
    verdachte_waarde = hoogste
else:
    verdachte_waarde = laagste

# analyse met afgeronde waarden op vier decimalen,
aantal_keer_std = round(afwijking_max / standaardafwijking, 4)
kans_kleiner = round(1 - norm.cdf(aantal_keer_std), 4)
kans_buiten = 2 * kans_kleiner
verwachtingswaarde = round(aantal_waarden * kans_buiten, 2)

# verwerpen of niet
if verwachtingswaarde < 0.5:
    wel_niet = "WEL"
else:
    wel_niet = "NIET"

print(waarden)
print("waarde " + str(verdachte_waarde) + " mag " + wel_niet + "
      worden verworpen")

```


Getallen

91. Zwevendekommagetallen

Dit is de langste hoofdstuktitel die uit maar één woord bestaat. En het is nog niet eens zeker of het de juiste vertaling van *floating point numbers* is. Het Engelse *floating* kan inderdaad ‘zweven’ betekenen, maar ook ‘drijven’.

Wikipedia gebruikt beide termen, en zelfs nog een derde.

Een zwevendekommagetal of drijvendekommagetal, verouderd ook vlotendekommagetal (Engels: floating-point number) is een gegevenstype dat een vaste geheugenruimte beslaat en een grote variëteit aan getallen kan bevatten, van zeer kleine tot zeer grote. Hoewel zwevendekommagetallen strikt genomen rationale getallen zijn, worden ze meestal gebruikt als benadering voor reële getallen. De relatieve nauwkeurigheid waarin getallen door zwevendekommagetallen worden gerepresenteerd, is min of meer gelijk over het gehele bereik. Het is de digitale versie van de wetenschappelijke notatie.

Als je de woorden onder elkaar zet, krijg je een grappig effect.

zwevendekommagetallen
drijvendekommagetallen

De eerste lijkt langer te zijn dan de tweede... terwijl het aantal letters juist eentje minder is: 21 in plaats van 22. Dat komt doordat ‘zwe’ (drie letters) breder is dan ‘drij’ (vier letters). De achttien letters daarna zijn hetzelfde.

In het geheugen van een computer wordt een zwevendekommagetal weergegeven in drie ‘delen’.

- *s*: een bit dat het teken aangeeft
- *e*: een bitreeks die een macht van 2 aangeeft
- *m*: een bitreeks die een geheel getal aangeeft

Een teken (*sign*, *s*), exponent (*e*) en mantisse (*m*).

Het woord ‘mantisse’ is ook een raar ding. Je komt het alleen hier tegen en in een paar verschillende betekenissen in de wiskunde. Het deel na de komma van een logaritmisch getal heet óók mantisse.

De representatie van een zwevendekommagetal ligt sinds 1985 vast in IEEE 754. De verdeling van de beschikbare bits verschilt bij de enkeleprecisievariant (*single*) van 32 bits en de dubbeleprecisievariant (*double*) van 64 bits.

- 32 bits: 1 voor *s*, 8 voor *e* en 23 voor *m*
- 64 bits: 1 voor *s*, 11 voor *e* en 52 voor *m*

Dat ene bit voor het teken is eenvoudig: een 1 is negatief, een 0 is positief.

Het binaire getal dat de exponent aanduidt, wordt voor de helft gebruikt voor negatieve getallen. Je moet daarom – we gaan nu voor het gemak even uit van 32-bits getallen – 127 van de decimale waarde aftrekken.

Bij de mantisse is het even opletten. Het binaire getal is namelijk *het deel na de komma*. Vóór de komma staat altijd een 1. En die 1 wordt niet gecodeerd in die 23 of 52 bits. Dat zou zonde van de ruimte zijn: hij is toch altijd 1.

Een voorbeeld: het getal 01000100110011010100000000000000

deel je op in die s, e en m: 0 10001001 100110101000000000000000

Die eerste 0 is een plus. De exponent rekenen we om naar decimaal: 137. Van die decimale 137 moeten we 127 aftrekken en dan houden we dus 10 over.

Die 100110101000000000000000 is het deel achter de komma, en staat dus voor 1,100110101. Alle laatste nullen zijn weggelaten, want die voegen niets toe. Dit binaire getal moet nog worden omgerekend naar decimaal: 1,603515625.

Dat omrekenen kun je via een *online tool* doen, maar ook “met de hand”. Het getal is $2^0 + 2^{-1} + 2^{-4} + 2^{-5} + 2^{-7} + 2^{-9} =$

$$1 + 0,5 + 0,0625 + 0,03125 + 0,0078125 + 0,001953125 = 1,603515625$$

In *Excel*:

1	0	1	1
1	-1	0,5	0,5
0	-2	0,25	0
0	-3	0,125	0
1	-4	0,0625	0,0625
1	-5	0,03125	0,03125
0	-6	0,015625	0
1	-7	0,0078125	0,0078125
0	-8	0,00390625	0
1	-9	0,001953125	0,001953125
			1,603515625

De waarde van het (positieve) getal kun je nu uitrekenen door de mantisse te vermenigvuldigen met 2 tot de macht die is uitgedrukt in de exponent.

$$1,603515625 \times 2^{10} = 1,603515625 \times 1024 = 1642$$

Nu andersom. We willen het getal pi opslaan als zwevendekommagetal. Pi is 3,141592653589793238462643383279502884197169399375105820974944592307.

Ongeveer dan. Het echte aantal decimalen is oneindig. Hoeveel kunnen we er kwijt? We hebben 23 bits voor de mantisse en $2^{-23} \approx 0,000000119$. De vijftien decimalen van *Excel* gaan we heus niet halen met 32 bits. En als we opslaan in 64 bits, houdt het ook een keer op. Dan zijn er 52 bits voor de mantisse en $2^{-52} \approx 2,22 \cdot 10^{-16}$. Dat komt in de buurt van *Excel*, maar is maar een kwart van wat hierboven staat op die ene regel. Dat lange getal, namelijk 3,141592653589793238462643383279502884197169399375105820974944592307.

De helft daarvan is (afgerond op diezelfde 66 decimalen):

1,570796326794896619231321691639751442098584699687552910487472296154.

We kunnen pi nu dus schrijven in de vorm $1, \dots \times 2^1$: het product van een gebroken getal met een 1 voor de komma en een macht van twee. Dit getal gaan we nu binair schrijven. Niet exact, maar in 66 decimalen, één regel vol.

1,100100100001111110110101010001000100001011010001100001000110100110.

De mantisse bestaat uit de eerste 52 bits na de komma.

$m = 1001001000011111101101010100010001000010110100011000$

We hadden al $s = 0$ (positief) en $e = 10000000000$ ($2^1 \rightarrow 1 + 1023$).

Achter elkaar gezet:

0 10000000000 1001001000011111101101010100010001000010110100011000

Hoe gaan we nu controleren of dit klopt? Tsja... Dat valt nog niet mee. Misschien moeten we het maar terugschroeven naar 32 bits. Daarna moet de exponent in 8 bits worden gecodeerd en de mantisse in 23 bits. Dan kun je met een *online tool* controleren of er hetzelfde uitkomt.

[IEEE-754 Floating Point Converter \(h-schmidt.net\)](http://h-schmidt.net/IEEE-754 Floating Point Converter)

0 10000000 10010010000111111011011

Joepie! Dat is hetzelfde.

92. De IEEE 754 Converter

Bestaat er dan geen *online tool* die voor ons uitzoekt hoe het zit?

Wat denk je nou zelf? Altijd als er in een boek of tijdens een les zo'n vraag wordt gesteld, is het antwoord "ja". Anders beginnen ze er niet over. Trouwens, als je opgelet had tijdens het lezen van de vorige pagina, had je het al gezien.

IEEE 754 Converter (JavaScript), V0.22

	Sign	Exponent	Mantissa
Value:	+1	2^1	1.5707963705062866
Encoded as:	0	128	4788187
Binary:	☐	☒ ☐ ☐ ☐ ☐ ☐ ☐ ☐ ☐	☒ ☐ ☐ ☒ ☐ ☐ ☒ ☐ ☐ ☐ ☐ ☒ ☒ ☒ ☒ ☒ ☒ ☐ ☒ ☒ ☐ ☒ ☒
You entered	3.14159265358979323846264338327950288420		
Value actually stored in float:	3.1415927410125732421875		
Error due to conversion:	8.742278000372485661672049712E-8		
Binary Representation	01000000010010010000111111011011		
Hexadecimal Representation	0x40490fdb		

De waarden verschillen dus onderling, en niet eens zo weinig. Het *online tool* kon ‘slechts’ 38 decimalen opslaan en we moesten het daarom afronden tot 3,14159265358979323846264338327950288420. Het getal dat uiteindelijk in 32 bits te coderen was, is 3,1415927410125732421875. Dat zijn 22 decimalen, maar alleen de eerste zes zijn exact kloppend.

Zes... Is dat veel of weinig? Als je 22 decimalen hebt en alleen de eerste zes kloppen, dan klinkt dat wat karig. Maar het is juist riant. Tien tot de macht zes is een miljoen, dus je kunt op een miljoenste nauwkeurig dat getal opslaan. Meestal vermenigvuldig je pi met iets waar meer onzekerheid op zit.

Moderne compilers werken met acht bytes voor een zwevendekommagetal. Daarmee heb je dus 52 bits om de getalswaarde vast te leggen en is de kleinste stapgrootte dus $2^{-52} \approx 2,2 \cdot 10^{-16}$. Dan zou je dus ongeveer vijftien decimalen aan moeten kunnen.

Onderstaand C programma laat zien dat dit aardig lukt.

```

int main(void)
{
    double test_pi =
        3.141592653589793238462643383279502884197169399375105820975;
    printf("%.20f\n", test_pi);
    printf("%d\n", sizeof(double));
    return 0;
}

>> 3.14159265358979310000
>> 8

```

Hetzelfde trucje in *Python* gaat ook goed.

```

test_pi = 3.141592653589793238462643383279502884197169399375105820975
print(test_pi)

>> 3.141592653589793

```

Er bestaan processoren met 128 bits. Als die alle 64 extra bits die ze hebben zouden gebruiken om een zwevendekommagetal nauwkeurig vast te leggen, zou je stapjes van 2^{-112} kunnen maken. Dan praat je over $1,9 \cdot 10^{-34}$. Een stuk aardiger al, maar nog steeds niet die 66 decimalen die op een regel passen.

Met 66 decimalen en één cijfer voor de komma heb je 67 cijfers. Daarmee kun je 10^{67} getallen opslaan, en de $^2\log$ van 10^{67} is ongeveer 222,69. Ervan uitgaande dat het aantal bytes een macht is van 2, hebben we 256 bits nodig. Niet overdreven, maar rekenen in 32 bytes vraagt om krachtige processoren.

Hoe kunnen we dat nou door 2 delen? Welke rekenmachine kan dat aan?

Er zijn websites die het wél kunnen. Op berekening.net staat er eentje die voor onze 66 decimalen zijn hand niet omdraait. Maar het is natuurlijk veel leuker om te doen alsof dat soort hulpmiddelen niet bestaat. <https://berekening.net/rekenmachine-voor-grote-getallen>

We kunnen dit probleem oplossen door het getal met 66 decimalen op te delen in brokjes van maximaal zes cijfers. Als we ervoor zorgen dat elk van die groepjes een even getal is, dan kun je ze één voor één delen door 2. Het lukt trouwens niet helemaal: in de reeks 939937510 zitten acht oneven cijfers achter elkaar.

93. Decimalen in Python en C

Een klassiek voorbeeld om te laten zien hoe het werken met *floats* mis kan gaan, is het optellen van 0,1 en 0,2.

Neem onderstaand Pythonprogramma.

```
a = 0.1
b = 0.2
c = a + b

print(a)
print(b)
print(c)

if 10 * a == 1:
    print("gelijk")
else:
    print("ongelijk")
```

Als je niet beter wist, zou je denken dat hier de getallen 0,1 en 0,2 en 0,3 worden afgedrukt, gevolgd door de conclusie dat tien keer een tiende gelijk is aan één.

Als je het ietsje beter weet en het vorige hoofdstuk begrepen en onthouden hebt, denk je dat die a niet precies 0,1 kan zijn en b niet precies 0,2 en c dus ook niet precies 0,3. En waar je vergif op zou innemen, is dat tien keer net-niet-een-tiende net niet gelijk is aan één.

Volgens ons hulpje [IEEE-754 Floating Point Converter \(h-schmidt.net\)](http://h-schmidt.net) zijn dit de “echte” getallen.

0,1	0,100000001490116119384765625
0,2	0,20000000298023223876953125
0,3	0,300000011920928955078125

Natuurlijk is het even schrikken als je ziet hoeveel cijfers er staan die daar helemaal niet horen. Maar kijk eens wat een lange lijst met nullen na die komma. Zo gek is dit nog niet. Het is alleen niet exact.

Nu de uitvoer van het programma.

```
0.1
0.2
0.30000000000000004
gelijk
```

Huh? Hij geeft die waarden van 0,1 en 0,2 helemaal goed weer.

Niet helemaal natuurlijk: er staat een punt terwijl er een komma zou moeten staan. Dat kun je *Python* niet kwalijk nemen. De taal is bedacht door de in Den Haag geboren Hollander Guido van Rossum, maar die heeft dat op het Engels gebaseerd.

Dat tien keer een tiende gelijk is aan één, is op zich waar. Maar de uitkomst van $10 \times 0,100000001490116119384765625$ is $1,00000001490116119384765625$. En dat is net geen één. We hadden dus ongelijk met ons 'ongelijk'.

Hoe gaat dit in C?

```
int main()
{
    float a = 0.1;
    float b = 0.2;
    float c = a + b;

    printf("%f\n", a);
    printf("%f\n", b);
    printf("%f\n", c);

    if (10 * a == 1)
        printf("gelijk");
    else
        printf("ongelijk");

    return 0;
}
```

Dit zou fout moeten gaan, toch? Nee hoor. Kijk maar naar de uitvoer.

```
0.100000
0.200000
0.300000
gelijk
```

Omdat ik het niet vertrouw, voer ik het programma nog een keer uit, maar dan door meer decimalen te gebruiken en de waarde van tien keer een tiende af te drukken.

```
printf("%.20f\n", a);  
printf("%.20f\n", b);  
printf("%.20f\n", c);  
printf("%.20f\n", 10 * a);
```

En maar weer eens naar de uitvoer kijken.

```
0.10000000149011612000  
0.20000000298023224000  
0.300000001192092896000  
1.00000000000000000000  
gelijk
```

Nou ja, zeg! Is hier nog een peil op te trekken? Een touw aan vast te knopen?

Terug naar het hulpje [IEEE-754 Floating Point Converter \(h-schmidt.net\)](http://h-schmidt.net/IEEE-754%20Floating%20Point%20Converter):
wat verwachten we precies?

0,1	0,100000001490116119384765625
0,100000001490116119384765625	0,100000001490116119384765625
1,00000001490116119384765625	1

Tsja. Een computer heeft dus moeite om 0,1 goed op te slaan. Of beter gezegd: het lukt niet. Het getal 0,100000001490116119384765625 gaat probleemloos, dankzij het feit dat het exact te schrijven is als het wat lange $1,600000023841858 \times 2^{-4}$ ($= 1,600000023841858 / 16$). Maar als je precies het tienvoudige van “zijn” versie van 0,1 probeert op te slaan, wat dus gelijk is aan 1,00000001490116119384765625, dan lukt ’m dat ook niet en zit hij er een klein stukje naast door exact 1 te nemen.

Logisch: 1 is een macht van 2 en komt wel heel dicht in de buurt.

Het getal dat in C wordt afgedrukt (0,10000000149011612000) is overigens niet exact gelijk aan 0,100000001490116119384765625. De eerste zestien decimalen zijn gelijk, met de laatste elf gaat het mis. Waar dát nou weer in zit...

94.Exacte getallen in Python

Als je 1 deelt door 2, krijg je “een half” of $\frac{1}{2}$ of 0,5. Allemaal best, niks mis mee. Deel je 1 door 4, dan komt daar “een kwart” uit of $\frac{1}{4}$ of 0,25. Ook al geen vuiltje aan de lucht. Maar voor sommige getallen geldt dat als je 1 erdoor deelt, dat er dan rare waarden op het glazige display van je rekenmachine verschijnen. Vreemde, maar vooral ook lange getallen. Met belachelijk veel cijfers achter de komma. Meer dan je eigenlijk zou willen. Oneindig veel soms zelfs.

Hieronder staan de getallen 1 t/m 30 en wat je ervan krijgt als je het getal 1 erdoor deelt.

1/1	1	1/16	0,0625
1/2	0,5	1/17	0,0588235294117647058823529...
1/3	0,3333333333333333333333333333...	1/18	0,05555555555555555555555555...
1/4	0,25	1/19	0,0526315789473684210526315...
1/5	0,2	1/20	0,05
1/6	0,16666666666666666666666666...	1/21	0,0476190476190476190476190...
1/7	0,142857142857142857142857...	1/22	0,0454545454545454545454545...
1/8	0,125	1/23	0,0434782608695652173913043...
1/9	0,11111111111111111111111111...	1/24	0,04166666666666666666666666...
1/10	0,1	1/25	0,04
1/11	0,0909090909090909090909090...	1/26	0,0384615384615384615384615...
1/12	0,08333333333333333333333333...	1/27	0,0370370370370370370370370...
1/13	0,076923076923076923079231...	1/28	0,0357142857142857142857143...
1/14	0,071428571428571428574286...	1/29	0,0344827586206896551724138...
1/15	0,06666666666666666666666666...	1/30	0,03333333333333333333333333...

[illegible]

Soms is een repeterende breuk leuk. Neem $1/6$. Dat begint gewoon met 0,1 en daarna pas komen er oneindig veel zessen. Of $1/18$: dat zijn allemaal vijven, maar wel met eerst een 0 achter de komma: 0,055555555555555555555556. Die laatste 6 is een nep-zes, maar als je correct wilt afronden, moet het zo.

Mijn favoriete repeterende breuk is $\frac{1}{22}$: 0,0454545454545454545454545.
Hij is nog net even lolliger dan z'n vader, $\frac{1}{11}$: 0,0909090909090909090909.

Gek hè? Als je op een verjaardag naast een onbekende gast komt te zitten en je raakt in gesprek over van alles, weten mensen soms moeiteloos hun favoriete restaurant of vakantieland te noemen. Politieke partijen en voetbalclubs kunnen ook, maar liggen wat gevoeliger. Maar vraag je iemand naar zijn of haar favoriete repeterende breuk, dan volgt vaak een lang en glazig zwijgen.

Een zeer eervolle vermelding is er voor de briljante breuk $1/7$, ook bekend als $0,142857142857142857142857142857142857142857142857142857142857142857$. Die repeteert ook, maar dat valt niet meteen op, omdat het met een blok van zes cijfers gaat. Op je rekenmachine staat bijvoorbeeld $0,142857143$. Hoezo repeterend? In dit getal van tien cijfers komen maar liefst acht verschillende voor. En na die tweede “14” komt eerst een 2 en dan een drie. Je zou niet zomaar op het idee komen dat een afgeronde herhaling is.

Rekenkundig kun je een repeterende breuk noteren door het cijfer of het groepje cijfers dat zich oneindig lang herhaalt te voorzien van een streep erboven. Wat je ook wel ziet, is haken eromheen.

$$\frac{1}{7} = 0,\overline{142857} = 0,[142857] \approx 0,142857142857142857142857142857$$

Let op: ook al staat het laatste getal in 24 decimalen: het is *ongeveer* gelijk aan $1/7$. Het isgelijkteken gebruik je alleen als het *exact* gelijk is.

Bij binaire getallen heb je ook repeterende breuken. Ze treden alleen op bij andere getallen dan de decimale. Dat geval van 0,1 dat we al bekeken hadden, is een mooi voorbeeld daarvan. Het decimale getal 0,1 schrijf je binair als:

[illegible]

We hebben bij het rekenen met exacte niet-gehele getallen eigenlijk met twee problemen te maken. Het gaat alleen maar goed als het getallen zijn die je zowel decimaal als binair kunt schrijven in een eindig aantal decimalen dat past binnen het aantal bits dat je hebt om het getal op te slaan.

Python heeft twee handige ‘dingen’ om onder deze problemen uit te komen, namelijk `Decimal` en `Fraction`.

Met `Decimal` kun je een getal dat als decimaal getal in een eindig aantal cijfers te schrijven is exact weergeven zonder dat je tegen de problemen aanloopt die met *floating point* representatie te maken hebben.

Met `Fraction` kun je een getal als breuk schrijven, dus in de vorm van een (gehele) teller en noemer.

Oefeningen

1. Schrijf een Pythonprogramma waarmee je een correcte optelling van de getallen 0,1 en 0,2 uitvoert.
2. Reken uit je hoofd uit hoeveel $2/3 + 3/7$ is.
 - a. Geef je antwoord in de vorm van een breuk.
 - b. Controleer je antwoord met een rekenmachine.
 - c. Programmeer deze som in *Python*.

Grafieken

95. Lijnen door punten

Als je één punt hebt, hoeveel lijnen kun je daar dan doorheen trekken?

Oneindig veel!

Als het er *oneindig veel* zijn, betekent dat dan ook dat *alle* lijnen door dat punt heen gaan? Nee. Oneindig veel betekent niet allemaal. Voor elke lijn die wel door het punt gaat, kun je oneindig veel lijnen bedenken die er niet doorheen gaan.

Het begrip *oneindig* is lastig, maar ga je beter bevatten als je het verhaal kent van die beheerder van een oneindig groot hotel. Alle kamers van het hotel (oneindig veel) zijn bezet: er zijn dus oneindig veel gasten. Nu komt er iemand binnen die vraagt of er nog een kamer vrij is. Nee, natuurlijk. Maar de beheerder bedenkt een list. Hij boekt de gasten die verblijven in kamer 1 om naar kamer 2. Daardoor komt kamer 1 vrij! Maar ja... waar laat je dan de gasten uit kamer 2? Gewoon, in kamer 3. Want degenen die daar slapen, kunnen óók een kamer opschuiven, naar kamer 4. En op dezelfde manier kunnen de mensen uit kamer 4 naar 5 en die uit N naar $N + 1$. Alle gasten (oneindig veel dus) schuiven één kamer op. Dat kan zonder problemen, want het aantal kamers eindigt niet. Na elke kamer komt er weer een kamer. En daarna weer eentje.

Duurt dat opschuiven dan niet oneindig lang? Nee hoor. Het gaat wel oneindig door, maar het duurt niet lang. Alle kamerwisselingen vinden tegelijkertijd plaats.

Er kunnen oneindig veel lijnen door één punt. En je kunt al die lijnen wiskundig nog beschrijven ook.

Voor elke rechte lijn geldt de volgende vergelijking.

$$y = ax + b$$

Voor alle punten (x_i, y_i) geldt dat er voor elke mogelijke waarde van a (en dat zijn er oneindig veel) een $b = y_i - ax_i$ bestaat waardoor die lijn door het punt (x_i, y_i) gaat.

Als je twee punten hebt, hoeveel lijnen kun je daar dan doorheen trekken? Eentje! Altijd precies één, tenzij de punten samenvallen. Dan kunnen er oneindig veel punten door, omdat je eigenlijk maar één punt hebt. Er is één lijn

mogelijk door twee niet-samenvallende punten (x_i, y_i) en (x_j, y_j) .
De richtingscoëfficiënt a ligt vast door die twee punten.

$$a = \frac{y_j - y_i}{x_j - x_i}$$

De waarde van b ligt ook vast.

$$b = y_i - ax_i$$

In plaats van het punt (x_i, y_i) mag je ook (x_j, y_j) invullen om b te bepalen. Er komt namelijk toch hetzelfde uit.

Als je drie punten hebt, hoeveel lijnen kun je daar dan doorheen trekken?

Nul of één!

Als dat derde punt exact op de lijn ligt die door de eerste twee gaat, dan past er precies één lijn door alle drie. (Er zijn trouwens oneindig veel punten op die lijn, maar er zijn nog veel meer punten die er *niet* op liggen.) Als dat derde punt net niet of helemaal niet op de lijn door de eerste twee ligt, dan gaat er geen lijn meer door de drie punten.

We hebben nu dus dit.

- één punt $\rightarrow$ oneindig veel lijnen
- twee punten $\rightarrow$ één lijn
- drie punten $\rightarrow$ nul of één lijn

Voor vier en meer punten geldt hetzelfde als voor drie. Ze kunnen allemaal op één lijn liggen. (Naarmate je meer punten hebt is de ‘kans’ dat dat zo is kleiner.)

Over de situatie met nul punten valt er ook iets te op te merken. Hoeveel lijnen gaan er door nul punten heen? Nul, zou je zeggen. Aan de andere kant: geef mij een willekeurig aantal punten en er is een lijn te bedenken die er niet doorheen gaat. Misschien zelfs oneindig veel. En misschien heeft het geen zin om hierover na te denken.

Door drie punten kun je wel altijd een parabool tekenen. En door vier punten past hoe dan ook een derdegraads polynoom. En door N punten een polynoom met de graad $N - 1$.

Bij metingen hebben we vaak veel meer punten dan nodig om er een lijn doorheen te trekken. Wat je dan zoekt, is de ‘best mogelijke’ lijn. Maar wat is de best mogelijke? Daarover gaat het in het volgende hoofdstuk.

Nog even over die beheerder van dat hotel. Stopt er opeens een oneindig grote touringcar met oneindig veel mensen erin. Die willen allemaal een kamer hebben. Dat gaat natuurlijk nooit passen. Hoewel... Als de gasten uit kamer 1 nu eens naar kamer 2 gaan, en die uit kamer 2 niet naar 3, maar naar 4. Dan komen kamer 1 én 3 alvast vrij, want die uit kamer 3 gaan naar 6. De gasten uit 4 gaan naar 8, die uit 5 naar 10. En die uit N naar $2N$. Zo komen alle oneven kamers beschikbaar. En dat zijn er oneindig veel. Want voor elke oneven kamer geldt: twee kamers verderop is er weer eentje.

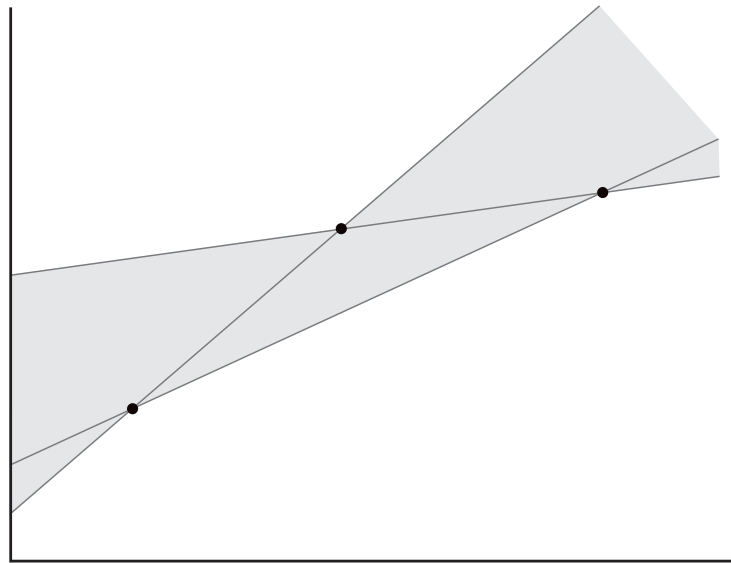
Gelukkig is de parkeerplaats oneindig groot en kan die bus gewoon achter het oneindig grote hotel staan.



96. Kleinste kwadraten

Je kunt dus door drie punten precies één lijn trekken, maar dan moeten ze wel precies in elkaars verlengde liggen. En meestal doen ze dat niet. Bijvoorbeeld in Figuur 96.1.

Als je toch een lijn wilt trekken, dan moet je ergens concessies doen. Ergens er-tussenin gaan zitten. Niet alles is even waarschijnlijk: normaal gesproken zul je toch iets kiezen in het grijze gebied.



Figuur 96.1 Het "grijze gebied" van lijnen door drie punten

Laten we een iets specifiekere geval kiezen, met coördinaten erbij. Welke lijn zou je trekken door de punten (2,2), (6,3) en (8,4)? Dat eerste punt ligt ergens links, de andere twee juist rechts. Als we er van de rechter twee eentje achterwege laten, vind je een oplossing.

$$y = ax + b$$

$$a = \frac{y_2 - y_1}{x_2 - x_1}$$

$$b = y_2 - ax_2 \quad \text{of} \quad b = y_1 - ax_1$$

Vul je hier de punten (6,3) en (2,2) in, dan vind je dat $a = 1/4$ en $b = 3/2$.

Vul je hier de punten (8,4) en (2,2) in, dan vind je dat $a = 1/3$ en $b = 4/3$.

Ergens daartussen moet het toch liggen, zou je zeggen. Maar wat is nou het beste? De beide a 's en b 's middelen? Dan moet je wel met breuken kunnen rekenen... Laten we alles in twaalfden doen.

Het gemiddelde van $1/4$ en $1/3$ is het gemiddelde van $3/12$ en $4/12$. Dat is $7/24$.

Het gemiddelde van $3/2$ en $4/3$ is het gemiddelde van $18/12$ en $16/12$. Dat is $17/12$.

Toch wordt zo'n 'middelste' lijn niet zomaar als de beste oplossing beschouwd. Meestal willen we de lijn zo trekken dat de verschillen tussen de lijn en de punten zo klein mogelijk zijn. En met 'zo klein mogelijk' bedoelen

we dan dat de kwadraten van de verschillen minimaal zijn. De waarden van a en b waarvoor dat geldt, volgen uit onderstaande formules.

$$a = \frac{N \sum xy - \sum x \sum y}{N \sum x^2 - (\sum x)^2}$$

$$b = \frac{\sum x^2 \sum y - \sum x \sum xy}{N \sum x^2 - (\sum x)^2}$$

Ja, daar word je een beetje duizelig van. Maar toch valt het wel mee. Tenminste: als je de moeite neemt om die termen en factoren even allemaal in een tabelletje te zetten.

x	y	xy	x^2
2	2	4	4
6	3	18	36
8	4	32	64
Σ	16	9	54
			104

Vier kolommen met daarin de drie punten en hun x en y en daarnaast het product xy en x^2 . De waarden in elk van die vier kolommen tel je op en je vindt vier getallen: 16, 9, 54 en 104. Met die vier getallen en de waarde voor het aantal punten N (in dit geval dus drie) kun je alles uitrekenen.

$$a = \frac{N \sum xy - \sum x \sum y}{N \sum x^2 - (\sum x)^2} = \frac{3 \times 54 - 16 \times 9}{3 \times 104 - 16^2} = \frac{162 - 144}{312 - 256} = \frac{18}{56}$$

$$b = \frac{\sum x^2 \sum y - \sum x \sum xy}{N \sum x^2 - (\sum x)^2} = \frac{104 \times 9 - 16 \times 54}{3 \times 104 - 16^2} = \frac{936 - 864}{312 - 256} = \frac{72}{56}$$

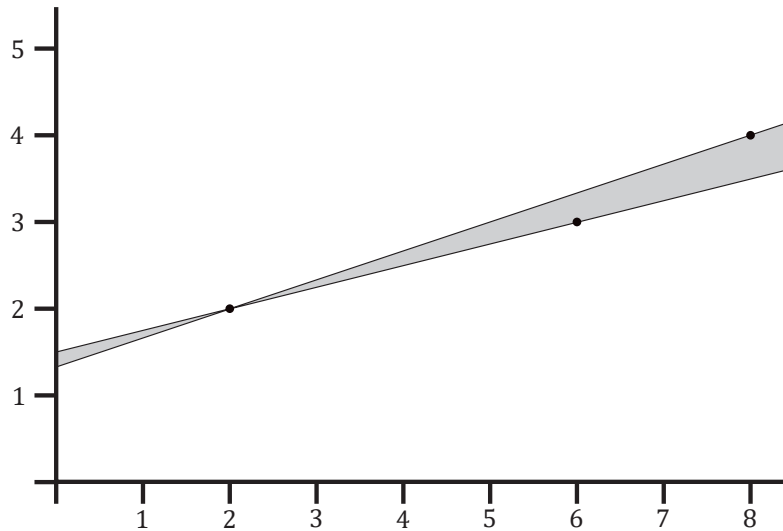
Het is eigenlijk niet moeilijk, maar het zijn stappen die je wel netjes één voor één moet kunnen uitvoeren. Zorgvuldig werken, daar gaat het om.

De waarden van a en b kun je ook als decimale getallen schrijven.

$$a = 0,3214\overline{285714}$$

$$b = 1,285714\overline{}$$

Het zijn exacte waarden, omdat de strepen boven de laatste cijfers aangeven dat het repeterende breuken zijn. Onafgerond dus. Exact.



Figuur 96.2 Twee mogelijke lijnen door drie punten

Maar moet het nou echt zo moeilijk? Nee.

Tik bij Google maar eens ‘least squares online’ in en je vindt heel wat websites die dit voor je kunnen uitrekenen. Als je daar de waarden invult, vind je hetzelfde. Meestal in een paar decimalen en/of significante cijfers:

$$a = 0,3214$$

$$b = 1,286$$

Van de repeterende breuk is niets meer te zien. Van de notatie in breuken al helemaal niet. Maar voor praktische situaties is het meestal goed genoeg.

In *Excel* is het ook al niet zo moeilijk. Zet de y-waarden in een kolom en x-waarden ernaast. (Let op: eerst de y en dan de x). De waarden van a en b vind je met deze formule. (De waarden van a en b worden in twee cellen getoond.)

`=LIJNSCH(A2:A4;B2:B4)`

En in *Python*? Is dat ingewikkeld? Nee. Je moet even de code weten.

```
from numpy.polynomial import polynomial as P

x = [2,6,8]
y = [2,3,4]

c, stats = P.polyfit(x,y,1,full=True)
print(c)
```

Deze code is trouwens voor het fitten van een willekeurige polynoom, dus niet per se een rechte lijn zoals hier. Door drie punten gaat altijd precies een parabool, dus je zou voor die 1 ook een 2 kunnen invullen.

97. Résidu

In het vorige hoofdstuk probeerden we ze goed mogelijk een rechte lijn te trekken door drie punten die niet op één lijn zaten met elkaar. We zagen dat er verschillende mogelijkheden waren, en dat er eentje is die we vaak als beste zien: de lijn waarvoor geldt dat de *som van de kwadraten van de residuen* zo klein mogelijk is.

Residuen? Ja. Zo hadden we ze nog niet genoemd, maar zo heten ze wel. Het verschil tussen de y-waarde van een punt en de y-waarde van de lijn die we gefit hebben bij de x-waarde van datzelfde punt, is een *residu*.

Ik vroeg me af wat de Nederlandse vertaling van residu is, omdat ik dacht dat het Frans was. Dus ik tikte ‘residu vertaling’ in bij Google en belandde (zoals ik gehoopt én verwacht had) bij Google Translate. En wat zei die? Die zei at de *Franse* vertaling van het *Nederlandse* woord ‘residu’ ‘résidu’ is. Even de puntjes op de i zetten: een streepje op de é. Het ‘echte’ Nederlandse woord is ‘overblijfsel’.

De residuen van drie punten reken je uit door de waarde op de gefitte rechte lijn $y_{fit} = 0,353214x + 1,286$ af te trekken van de ‘werkelijke’ y.

x	y	y_{fit}	residu	residu ²
2	2	1,9288	-0,0712	0,005069
6	3	3,2144	0,2144	0,045967
8	4	3,8572	-0,1428	0,020392

Laten we voor de gein onze ‘eigen’ rechte lijn hier eens mee vergelijken. We hadden zelf al een keer een rechte lijn bedacht die door het linker punt ging en zo’n beetje midden tussen die andere twee door. Dat paste ook wel aardig, toch? We kregen er exacte breuken uit.

En een exacte breuk is wel zo leuk.

Om het residu in elk punt te vinden, gaan we dus

$$y_{onze\ eigen(wijze)\ fit} = \frac{7}{12}x + \frac{24}{17}$$

aftrekken van de ‘werkelijke’ y.

x	y	$y_{o.e.w.f.}$	residu	residu ²
2	2	2	0	0
6	3	3,166667	0,166667	0,027778
8	4	3,75	-0,25	0,0625

Past die lijn van ons *beter* dan die van de kleinste kwadraten methode? De residuen zijn kleiner... Zowel in absolute waarden als in de kwadraten. Kijk maar in de tabel. Voor het punt (2,2) is 'ons' residu exact *nul*. Voor het punt (6,3) is het kwadraat ongeveer de helft. Alleen voor dat derde punt is hij iets groter. Het tweede past beter, het derde slechter, het eerste is *exact* goed. Een twee derde meerderheid van de stemmen zegt dat onze lijn beter is.

Maar ja, wanneer is iets *beter*? Als je naar de harde feiten kijkt en de som van de kwadraten optelt, is die van de door *Excel* en *Python* of wat dan ook gefitte lijn inderdaad beter. Tel ze maar bij elkaar op. De som van 'onze' lijn is afgerond 0,09. Die van 'hun' lijn is afgerond 0,07. Het scheelt ruim een kwart.

0,005069	0
0,045967	0,027778
0,020392	0,0625
0,071429	0,090278

Maar waarom kijken we eigenlijk steeds naar de *kwadraten* van de verschillen? Als ik wil weten of A dichterbij B ligt dan bij C, dan ligt het toch voor de hand om naar de *absolute* verschillen te kijken?

Arend haalt een 6,1 voor het tentamen en Bert een 9,6. Cor had een 5,4. Bij wie komt Arend het dichtst in de buurt? Bij Cor, want dat scheelt maar een halve punt en het verschil met Bert is drieënhalve punt. Dat is een factor zeven meer! Toch zitten Arend en Bert samen in dezelfde categorie, namelijk de club die de taart kan aansnijden. Die arme Cor is gezakt. Zo zie je hoe betrekkelijk getallen en cijfers zijn. De meest relevante vraag is in dit geval of het een voldoende was.

Als we de absolute waarden van de residuen nemen, krijgen we een andere vergelijking. Het scheelt heel weinig (zo'n 2,8%), maar 'onze' simpele fit is volgens de som van de absolute waarden iets kleiner dan de lijn die via de methode van de kleinste kwadraten gevonden is.

0,0712	0
0,2144	0,1667
0,1428	0,25
0,4284	0,4167

Zouden er nog andere manieren en criteria te bedenken zijn om te kijken naar een 'beste' lijn? Jazeker. Of het de beste is, is altijd nog een vraag, maar op zijn minst kun je naar een zinvolle lijn zoeken.

Wat te denken van de lijn die het grootste residu zo klein mogelijk maakt? Uitgaande van de keuze om de lijn links door dat ene punt te laten gaan en

rechts ergens tussen die andere twee door, kunnen we beginnen met vast te stellen dat het residu van het eerste punt per definitie gelijk is aan nul.

Voor een punt (x_i, y_i) geldt:

$$r_i = y_i - ax_i - b$$

Het kleinst mogelijke maximale verschil krijg je als de verschillen qua grootte aan elkaar gelijk zijn. Het ene positief en het andere negatief. In alle andere gevallen zijn de verschillen namelijk groter: als het ene residu afneemt, neemt het andere toe.

$$y_2 - ax_2 - b = -(y_3 - ax_3 - b)$$

$$y_2 + y_3 - a(x_2 + x_3) = 2b$$

$$b = \frac{y_2 + y_3 - a(x_2 + x_3)}{2} = \frac{3 + 4 - a(6 + 8)}{2} = 3,5 - 7a$$

We hebben nu één vergelijking met twee onbekenden, en hebben er dus nog eentje nodig. Die volgt uit dat andere punt (helemaal links) waarvoor geldt dat het residu nul is.

$$0 = y_1 - ax_1 - b$$

$$b = y_1 - ax_1 = 2 - 2a$$

Deze vergelijking kunnen we combineren met de vorige.

$$2 - 2a = 3,5 - 7a$$

$$5a = 1,5$$

$$a = 0,3$$

$$b = 2 - 2a = 1,4$$

Een andere lijn dus, die volgens een ander criterium 'beter' is.

98. Raad het volgende getal

Hughes & Hase hebben een alleraardigst raadseltje, dat aangeeft dat het fitten van een hoge orde functie op een eindig aantal getallen altijd wel een perfecte pasvorm oplevert, maar niet altijd zinnig is.

Voordat ze hun raadseltje opgeven, staan ze eerst stil bij het begrip *vrijheidsgraad*. Als je N onafhankelijke datapunten hebt en je fit een functie met P parameters, dan is het aantal vrijheidsgraden v (*degrees of freedom*)

$$v = N - P.$$

Ten onrechte denken sommige mensen dat het aantal vrijheidsgraden bij voorkeur zo klein mogelijk is. Dan past de fit beter. Natuurkundig gezien wil je vaak juist zo veel mogelijk vrijheidsgraden hebben: een fit met zo weinig mogelijk parameters, die toch de data goed beschrijven.

Nu het raadseltje: wat is het volgende getal in de reeks 1, 2, 4, 6, 10?

Je bedenktijd is voorbij! Het juiste antwoord is: 12. Waarom? Omdat de reeks gevormd wordt uit “de priemgetallen min 1”. De priemgetallen zijn 2, 3, 5, 7, 11, 13, 17, 19 etc. Als je daar 1 van aftrekt, vind je 1, 2, 4, 6, 10, 12, 16, 18 etc.

Het lukt niet om een lineaire functie door deze punten trekken. Kwadratisch gaat al iets beter en een derde macht begint ergens op te lijken. Omdat het vijf punten zijn, past een vierdegraads polynoom exact op de data. Een rechte lijn gaat immers precies door twee punten, een parabool door drie...

Deze functie past volgens Hughes & Hase.

$$f(x) = 5 - 8,5833x + 5,875x^2 - 1,4167x^3 + 0,125x^4$$

Hier kun je $x = 6$ invullen en dan komt er 21 uit. Niet het juiste antwoord. Toch wel! Dit antwoord is *ook* juist.

Oefeningen

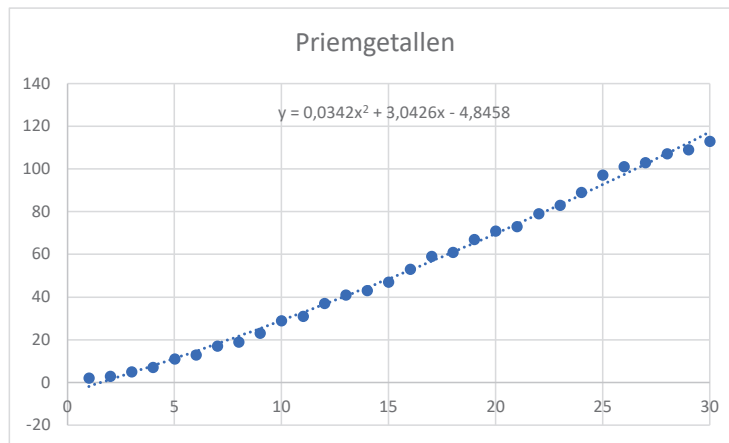
1. Controleer of bovenstaande functie klopt, op de volgende manieren.
 - a. een oplossing zoeken via een online tool
 - b. een programma in *Python*
 - c. een grafiek tekenen met *Excel* en dan een polynoom fitten
 - d. de door Hughes & Hase gevonden oplossing controleren door die in te vullen voor $f(0)$, $f(1)$, $f(2)$, $f(3)$, $f(4)$, en $f(5)$
2. Bedenk een andere juiste oplossing voor het raadseltje.

99. Past het model?

Het raadseltje in het vorige hoofdstuk kun je leuk vinden (of niet), maar er zit natuurlijk meer achter dan luim of lol. We hadden de reeks 1, 2, 4, 6, en 10 en bleken een functie te kunnen vinden die met de getallen 5, $-8,5833$, $5,875$, $-1,4167$ en $0,125$ die reeks bijna helemaal goed beschreef.

Bijna helemaal goed, maar vraag niet hoe. Deze beschrijving wordt gegeven door vijf getallen (met samen zeventien significante cijfers), voor een reeks van vijf andere getallen (met samen zes significante cijfers)...

Als je de 5 in bovenstaande reeks parameters vervangt door een 6, dan zou je niet de priemgetallen -1 hebben, maar de priemgetallen zelf. In een grafiek zien de eerste dertig priemgetallen er als volgt uit.



Nee, ze liggen niet op een rechte lijn en er is geen simpele formule om de volgende waarde exact te voorspellen. Maar zelfs een simpele parabool kan zodanig gefit worden dat alle dertig verschillen onder de 4,5 blijven. Een vierdegraadsvergelijking om het een beetje beter te doen, is overdreven.

Er is geen polynoom die deze reeks beschrijft. Althans: niet eentje die ook past op het volgende getal, als je geen orde gebruikt die net zo hoog is als (of hoger dan) het aantal getallen.

Wat nu als je een polynoom probeert te fitten op data waar wél een wiskundig verband achter zit? Huges & Hase gebruiken – na dat raadseltje van het vorige hoofdstuk – de wiskundige beschrijving die het verband aangeeft tussen slingerijd, slingerlengte én beginhoek.

$$T = 2\pi \sqrt{\frac{l}{g}} \left(1 + \frac{1}{16} \theta^2\right)$$

Dit verband is een vereenvoudiging van de eerste-orde correctie op de formule die je meestal ziet en die ‘geldig’ is voor kleine hoeken.

$$T = 2\pi \sqrt{\frac{l}{g}}$$

De ‘echte’ analytische oplossing bevat oneindig veel termen.

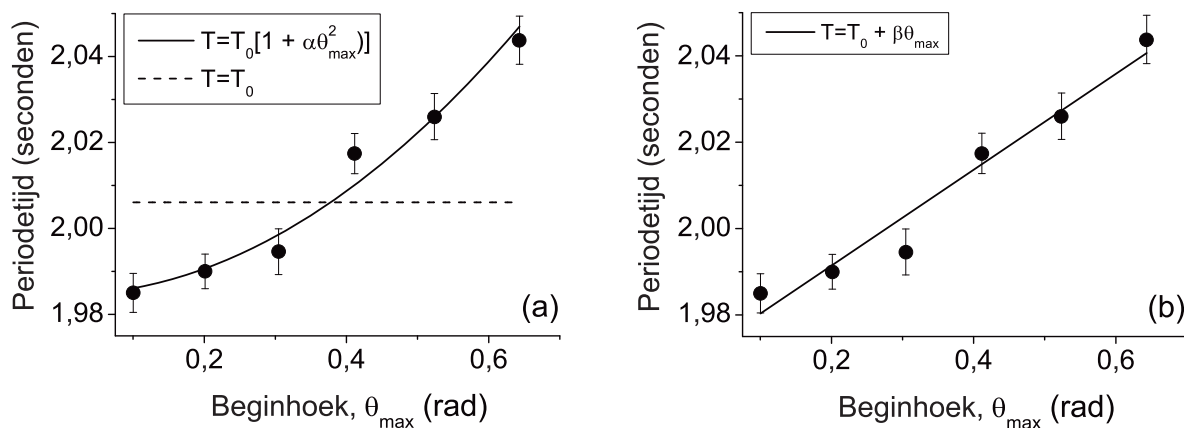
$$T = 2\pi \sqrt{\frac{l}{g}} \left(1 + \frac{1}{16} \theta^2 + \frac{11}{3072} \theta^4 + \frac{173}{737280} \theta^6 + \frac{22931}{132120570} \theta^8 + \dots \right)$$

Je bent al aardig op weg als je alleen de eerste twee neemt.

$$T = 2\pi \sqrt{\frac{l}{g}} \left(1 + \frac{1}{16} \theta^2 \right)$$

Die formule met die ene term met θ^2 erin is dus kwadratisch afhankelijk van de beginhoek als je de hogere-orde termen negeert. Dat kun je theoretisch afleiden, maar blijkt het ook uit metingen?

Huges & Hase geven een voorbeeld waarbij ze niet de meetdata zelf geven, maar wel een grafiek met meetresultaten. Grafisch ziet het er als volgt uit.



In beide plaatjes staan dezelfde meetpunten. Links is een poging gedaan om een parabool door de foutgebieden te trekken. Dat lukt net niet. Rechts is geprobeerd om dat met een rechte lijn te doen. Ook dat slaagt net niet helemaal. Vijf van de zes foutgebieden omvatten de lijn, eentje niet.

En dan is er ook nog die stippellijn (gestreepte lijn) die uitgaat van een constante waarde. Het lukt je niet om een horizontale lijn te trekken die door meer dan twee foutgebieden gaat. Belangrijker dan dat: je ‘ziet meteen’ dat

de zes meetpunten een oplopende volgorde hebben. Voor alle zes punten geldt: als de hoek groter is, dan is de periodetijd ook groter.

Hughes & Hase hebben de Chi-kwadraattoets gebruikt om te toetsen hoe het zit. De wiskunde die daarachter zit, komt in *Meten & Weten* niet aan de orde. In het kort: een Chi-kwadraattoets levert een getal op dat in de buurt van 1 moet liggen. Als het 2 of hoger is (bij dit aantal vrijheidsgraden), moet je de zogenaamde nulhypothese verwerpen.

Drie nulhypotheses worden er door Hughes & Hase getoetst, namelijk de hypothesen dat de data beschreven worden door de volgende drie modellen.

$$T = T_0$$

$$T = T_0(1 + a\theta)$$

$$T = T_0(1 + b\theta^2)$$

De *waarden* (die dus rond het getal 1 moeten liggen en in ieder geval kleiner moeten zijn dan 2) voor deze drie modellen zijn achtereenvolgens 21,4, 0,9 en 1,1. Op basis daarvan wordt het eerste model duidelijk verworpen en zijn de andere twee heel goed mogelijk.

Een model dat buiten beschouwing wordt gelaten, is het volgende.

$$T = T_0(1 + c\theta + d\theta^2)$$

Terwijl je met zekerheid kunt zeggen dat dit model óók past: als c of d gelijk is aan 0, dan hebben we namelijk die andere twee, waarvan we al weten dat ze pasten.

Oefening

1. Bereken of beredeneer dat het verdedigbaar is om de termen die van een hogere orde zijn dan de kwadratische term te verwaarlozen zijn.

Rekenen

100. De taylorreeks van een sinus

De sinus van een hoek is gedefinieerd als de lengte van de overstaande rechtehoekszijde gedeeld door de lengte van de schuine zijde. Voor een heel kleine hoek is dat een heel klein getal: de overstaande zijde is namelijk heel kort.

Wiskundig kan worden aangetoond dat je een sinus kunt schrijven als een taylorreeks.

$$\sin(x) = \sum_{n=0}^{\infty} \frac{(-1)^n}{(2n+1)!} x^{2n+1}$$

Wiskundigen worden blij van dit soort formules. Ze kloppen exact en zijn compact weergegeven, met zo min mogelijk concrete getallen en cijfers erin. De 0 die er staat, geeft aan dat je alle mogelijke getallen uit de verzameling $\mathbb{N}$ (natuurlijke getallen) afloopt. De enen en tweeën staan voor verhogingen en verdubbelingen, maar je komt in deze formules geen vijven en zevens tegen, laat staan getallen als 120 (de faculteit van 5).

Niet-wiskundigen vinden het vaak prettiger lezen als je dat somteken weglaat en juist wél concrete gevallen invult, gevolgd door drie puntjes.

$$\sin(x) = x - \frac{1}{3!}x^3 + \frac{1}{5!}x^5 - \frac{1}{7!}x^7 + \dots$$

Dit is geen exacte formule meer, maar je snapt wel hoe het verder gaat. Minder overzichtelijk wordt het als je de faculteiten uitrekent en opschrijft, maar het geeft je wel meer gevoel voor groottes.

$$\sin(x) = x - \frac{1}{6}x^3 + \frac{1}{120}x^5 - \frac{1}{5040}x^7 + \frac{1}{362880}x^9 - \dots$$

Aan 362 880 zie je niet meteen dat het de faculteit van 9 is en 39 916 800 herken je ook niet meteen als de faculteit van 11. Hoewel... die 00 aan het einde geeft aan dat het veelvoud van 100 is en dat komt door dat 2×5 en 10 er allebei in zitten.

Maar wat aan 39 916 800 weer wél handig is, is dat je in één oogopslag ziet dat het veel groter is dan 120. Terwijl voor je gevoel (als je niet veel met getallen, rekenen en wiskunde hebt) er niet zo'n heel groot verschil is tussen 5! en 9! Het lijkt iets niet al te groots zijn als 4!, maar het is altijd nog wel een factor $9 \times 8 \times 7 \times 6 = 3024$.

Nu nog die x . Als we het over een hoek hebben, moet dat een getal zijn tussen 0 en 2π . Een rechte hoek is π radialen, scherpe hoeken zijn kleiner dan π .

Wat nu als het een klein getal is? Wat is klein? Laten we zeggen: 0,1 radialen, dat is nog geen 6°, ongeveer 5,73° om niet precies te zijn. Wat gebeurt er in die formule met die (oneven) machten?

$$\sin(0,1) = 0,1 - \frac{0,001}{6} + \frac{0,00001}{120} - \frac{0,0000001}{5040} + \frac{0,000000001}{362880} - \dots$$

Nog verder uitschrijven levert:

$$0,1 - 0,0001667 + 0,00000008333 - 0,0000000001984 + 0,000000000...$$

Je ziet dan dat alleen de tweede term nog van enig belang is. Dat klopt ook met wat je ziet als je $\sin(0,1)$ intikt op je rekenmachine. Dat is namelijk 0,09983, nog geen 2‰ verschil met 0,1.

“Voor kleine hoeken is $\sin(x)$ ongeveer gelijk aan x ”, zeggen ze wel eens. Dit is waarom dat zo is.

Oefeningen

1. Bereken op je rekenmachine de sinus van 0,01 radialen.
Bereken hoeveel promille die uitkomst afwijkt van 0,01.
2. Bereken met een Taylorreeks de sinus van 0,001 radialen.
Bepaal in hoeveel significante cijfers je het antwoord moet geven om te laten zien dat het niet exact gelijk is aan 0,001.

101. Graden

Wat is het verschil tussen een rechte hoek en het kookpunt van water?

Als je zegt dat het tien graden is, dan zie je het onzinnige van de vraag niet in. Of juist wel. Een graad op je geodriehoek is iets heel anders dan een graad (Celsius of Fahrenheit) op de thermometer. Het is ook geen academische graad of een graat in je gebakken vis.

Dat een vissengraat geen graad is, zie je aan de spelling. Dat een graad in de temperatuur iets anders is dan een graad in een hoek, zie je aan de C of F die erachter staat. En nu we het er toch over hebben: het woord *graden* (of het teken $^{\circ}$) wordt niet gebruikt als we het over de eenheid kelvin hebben. Dat is pas in 1967 besloten, dus in oude boeken kan het ook anders staan. Eén kelvin (dat is dus 1 K) is gedefinieerd als het $1/273,16$ deel van het verschil tussen het absolute nulpunt en het *tripelpunt* van water. Het tripelpunt heeft niets met Belgisch bier te maken, maar is de temperatuur waarop een stof tegelijkertijd kan voorkomen als vaste stof, vloeistof en gas.

Dat tripelpunt van water ligt op $0,01^{\circ}\text{C}$: nét iets boven wat we in de volksmond het vriespunt noemen. Het goede nieuws is dat het verschil zo klein is dat het in praktijk geen betekenis heeft. Geen enkele weerman verwacht dat het morgen $18,1^{\circ}\text{C}$ wordt. Ze zeggen allemaal ‘achttien graden’, zonder iets achter de komma en zonder de toevoeging Celsius.

Laten we even de dingen op een rij zetten. Als het 20°C is, dan kun je dat noemen als

$$T = (20,0 \pm 0,5)^{\circ}\text{C}$$

Je geeft daarmee aan dat je het op hele graden Celsius nauwkeurig weet. Het is heus geen 21°C en ook geen 19°C , maar daartussenin. Let op dat je een spatie zet tussen het getal en de eenheid. Volgens de internationale ISO-normen (en de Nederlandse norm NEN EN ISO 80000-1:2013) hoort dat zo, maar taalkundigen raden soms aan om 20°C te schrijven. Bij de 90° van een rechte hoek noteer je het gradenteken trouwens wel achter het getal.

Als je dezelfde temperatuur zou willen schrijven in kelvin (wat natuurkundigen graag doen), dan zou je wat de notatie betreft misschien zoiets willen.

$$T = (293,16 \pm 0,05) \text{ K}$$

Uiteraard mag dat niet: het is wel leuk om te laten zien dat je weet dat je netjes gerekend hebt met 273,16 K en niet met 273,15 K, maar de oorspronkelijk

meting gaf al aan dat je onzekerheid groter was. Je moet dan ook afronden op één cijfer achter de komma.

$$T = (293,2 \pm 0,5) K$$

De eenheid is met een hoofdletter K, maar als je het als woord uitschrijft in de tekst is het juist met een kleine letter: *kelvin*. Nogmaals: geen *graden* kelvin, maar gewoon *kelvin*. Net als meter en volt en ampère: dat is ook met kleine letters en zonder graden.

Een graad Celsius is zo leuk omdat je de temperatuur uitdrukt in een soort percentage van een heel alledaagse schaal. Als 0 °C het vriespunt en 100 °C het kookpunt van water is, dan zijn de getallen daartussenin waarden die betekenis voor je hebben. Als je dan ook nog weet dat kamertemperatuur 20 °C is, je koelkast op 5 °C staat, dat je in water van 50 °C nog net je handen kunt houden zonder dat het pijn doet en lichaamstemperatuur iets minder dan 37 °C, dan krijg je een ‘gevoel’ bij die waarden.

Waarom rekenen natuurkundigen dan liever in kelvin? Omdat je dan een absolute schaal hebt die bij nul begint, en je dus ook kunt vermenigvuldigen en delen. Als het de ene dag 10 °C is en de andere dag 20 °C, dan is het niet ‘twee keer zo warm’. Maar als je in kelvin rekt, kun je wel over een verdubbeling van de temperatuur spreken bij het toepassen van de algemene gaswet.

$$pV = nRT$$

Als de temperatuur (in kelvin) 10% hoger wordt en het volume blijft gelijk, dan wordt de druk ook 10% hoger.

Oefeningen

1. Als het de ene dag 10 °C is en de andere dag 20 °C, dan is het niet ‘twee keer zo warm’.
Hoeveel procent warmer is het dan wel?
2. Jeroen Jansen wil graag een eigen temperatuurschaal bedenken. Hij twijfelt of hij de eenheden *graden Jansen* of *graden Jeroen* wil gaan noemen. In ieder geval gaat het genoteerd worden in °J. Het absolute nulpunt noemt hij 0 °J en het tripelpunt van water noemt hij 100 °J. Bereken hoeveel graden Jeroen (°J) overeenkomt met 20 °C.

102. Radialen

We rekenen makkelijker in kelvin dan in graden Celsius (of Fahrenheit of graden). En voor hoeken liever met radialen dan graden: 90° is een handig getal, want je kunt het delen door 2, 3, 5, 6, 9, 10, 15, 18, 30, 45 en 90. Dat zijn maar liefs elf getallen, terwijl je 100 maar door acht getallen kunt delen. (Uiteraard kun je ze allebei ook door 1 delen. Maar dan deel je niet.)

Het lijkt aardig om een rechte hoek te definiëren als 60° . Het getal 60 heeft namelijk ook veel delers: 2, 3, 4, 5, 6, 10, 12, 15, 20, 30, 60. Dat zijn er ook elf, terwijl het getal 60 zelf een kleiner is dan 90. Maar de keuze voor 90 is handiger, omdat de delers wat beter verspreid zijn. Kijk maar in de tabel hieronder.

deler	7	12	60	90	80	100
2	nee	ja	ja	ja	ja	ja
3	nee	ja	ja	ja	nee	nee
4	nee	ja	ja	nee	ja	ja
5	nee	nee	ja	ja	ja	ja
6	nee	ja	ja	ja	nee	nee
7	ja	nee	nee	nee	nee	nee
8	nee	nee	nee	nee	ja	nee
9	nee	nee	nee	ja	nee	nee
10	nee	nee	ja	ja	ja	ja

Maar waarom dan toch rekenen met radialen? Graden zijn toch veel makkelijker om je iets bij voor te stellen? Als 90° een rechte hoek is, dan is 81° bijna een rechte hoek: je zit maar 10% van de 90° af, daar kun je je meteen iets bij voorstellen. Maar wat zegt een hoek van 1,4 radialen? Dat moet je dan door π delen en met 180° vermenigvuldigen om te zien dat het tussen de 80° en 81° zit. Je kunt je daar minder goed iets bij voorstellen.

Toch doen wiskundigen het niet zomaar. Het leuke van radialen is dat je de oppervlakte van een taartpunt kunt berekenen door de hoek (in radialen) te vermenigvuldigen met de straal in het kwadraat en dat te delen door 2.

Hoe zit dat? Een hele taart is rond en heeft een oppervlakte van πr^2 . Die hele taart is een 'hoek' van 360° oftewel 2π radialen. Het deel dat je daarvan hebt (de hoek α) is $\alpha/2\pi$. De oppervlakte van een taartpunt is dus:

$$A = \frac{\alpha}{2\pi} \cdot \pi r^2 = \alpha \cdot \frac{r^2}{2}$$

103. 3x

“Er zijn er 3x meer dan gisteren. Gisteren waren er 37.
Hoeveel zijn er dan vandaag?”

Lastige vraag... En dat komt vooral doordat we niet weten wat er wordt bedoeld. Wat is die ‘3x’? Sommige lezers zullen het opvatten als ‘drie keer’. Maar dat staat er niet! Er staan een cijfer 3 en een letter x.

Een x is een letter. Een × is een vermenigvuldigingsteken.

Als je de tekst cursief afdrukt, is een x al helemaal geen keerteken meer. Een keerteken maak je met Alt + 0215. Of door het getal 00D7 in te tikken en daarna At + X. Makkelijker kunnen we het niet maken: vijf of zes keer een toets in drukken voor één keer ‘keer’...

Wiskundig klopt het niet als je een x gebruikt als keerteken. Als y tien keer de waarde van x heeft, dan moet je niet schrijven: $y = x x x$, ook al is het Romeinse cijfer X een 10. Correct is: $y = 10 \times x$.



In een schreefloze letter zoals *Arial* lijkt een x nog aardig op een ×. Bij een schreefletter als *Times New Roman* wordt dat al minder. Voor beide geldt dat het keerteken hoger op de regel staat.

Taalkundig klopt het ook niet. Het keerteken (of de x) is geen afkorting van ‘keer’. En getallen tot en met twintig (en tientallen, hondertallen...) schrijf je in principe uit in letters. *Onze Taal* geeft daar [regels](#) voor. Het Vlaamse *Team Taaladvies* ook, en die [regels](#) zijn hetzelfde. Ook de Nederlands-Belgische samenwerking *Taaladvies.net* heeft dezelfde [regels](#).

Die regels zijn geen officiële regels. Dat staat ook op de genoemde sites. Toch doe je er als taalgebruiker goed aan om je te conformeren aan wat gebruikelijk is. Het raadplegen van gezaghebbende bronnen is nuttig.

Logisch gezien is het ook nog lastig. Is *drie keer meer* méér dan *drie keer zoveel*? ‘Drie keer meer’ klinkt als ‘vier keer zo veel’... Is één keer meer hetzelfde als het dubbele? Ook daar schrijf *Onze Taal* [iets](#) over.

Wat bedoelt iemand eigenlijk met 3x? Drie keer de waarde van variabele x is een mogelijkheid. Maar x als willekeurig cijfer in een getal met twee cijfers kan ook. Dan betekent 3x dus één van de getallen 30 t/m 39.

Vroeger, toen de meeste mensen thuis nog geen computer hadden, werden brieven en verslagen – ook voor school – vaak nog gemaakt met een schrijfmachine. Een mechanisch apparaat met toetsen die een soort stempels aandreven en een inktlint. Sommige van die schrijfmachines hadden wel de cijfers van 2 t/m 9, maar geen 0 en 1. Wie een 0 of 1 nodig had, gebruikte dan maar de kleine letter l of de hoofdletter O.

10 x 10 = 100

In het lettertype *Urania Czech* hebben ze de kleine letter l en de 1 aan elkaar gelijk gemaakt. De 0 en de O lijken ook op elkaar. Er zit een klein hoogteverschil in, waarschijnlijk om de onvolmaaktheid van de schrijfmachine na te bootsen.

Het eerste getal 10 in de som hierboven is gemaakt met letters.
Het tweede getal 10 in de som hierboven is gemaakt met cijfers.

Op Nederlandse schrijfmachines had je ook vaak nog een toets voor de gulden en de digraaf ij: de f en de y.

f y

Eigenlijk is het vreemd dat sommige cijfers en letters zo veel op elkaar lijken. En dat woorden soms twee verschillende betekenissen kunnen hebben. Neem onderstaande Engelse zin.

III... I'll lie

De hoofdletter I is in *Arial* nog geen 8% dikker dan de kleine letter l. Staat hier nu “Ziek... Ik zal liggen”, of “Ziek... Ik zal liegen”?

Aan de andere kant lijken hoofdletters soms maar heel weinig op de kleine letter waar ze bij horen. In de woorden ‘GAT’ en ‘gat’ zie je dat de drie letters veel van elkaar verschillen.

104. Maak er maar geen punt van

Moet je in de uitkomst van 2 gedeeld door 5 een punt of een komma zetten?

Is het 0.4 of 0,4? Mijn rekenmachine zei 0.4...

Even voor alle duidelijkheid: 2 door 5 delen lukt me ook wel zonder rekenmachine. Die gebruikte ik alleen maar om te kijken of er een punt of een komma zou komen. Dat sommetje ging uit het blote hoofd.

Mijn rekenmachine is beter in rekenen dan in taal. Helaas heb ik geen taalmachine. Dan maar met Google aan de slag.

Scribbr: In het Nederlands gebruik je komma's als decimaalteken, terwijl je in het Engels een punt gebruikt.

Taaladvies.net: In Nederland en België wordt een komma gezet tussen 'hele getallen' en decimalen (die niet voor niets ook wel 'de cijfers achter de komma' genoemd worden).

Wikipedia: In het Nederlands gebruikt men een komma als decimaalteken, terwijl Engelstaligen een punt gebruiken.

Onze Taal: In een getal met een of meer decimale cijfers wordt in het Nederlands traditioneel een komma gebruikt: 36,8 °C, 1,98 m, 3,5 miljoen. De punt komt onder invloed van het Engels wel steeds meer voor.

Vlaanderen.be: Volgens het Bureau voor Normalisatie (NBN) en het Nederlands Normalisatie-instituut (NEN) gebruiken we in getallen een komma als decimaalteken en een spatie om de cijfers in groepjes van drie van elkaar te scheiden.

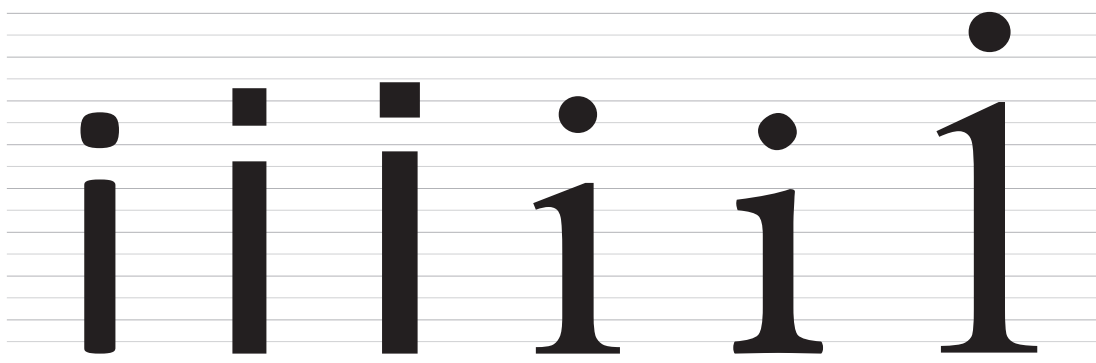
Alleen bij Onze Taal laten ze doorschemeren dat het misschien aan het veranderen is. Mensen hebben het ook over kids als ze kinderen bedoelen.

De Nederlandstalige versie van Excel en de Rekenmachine van Windows vertonen ook een komma in getallen.

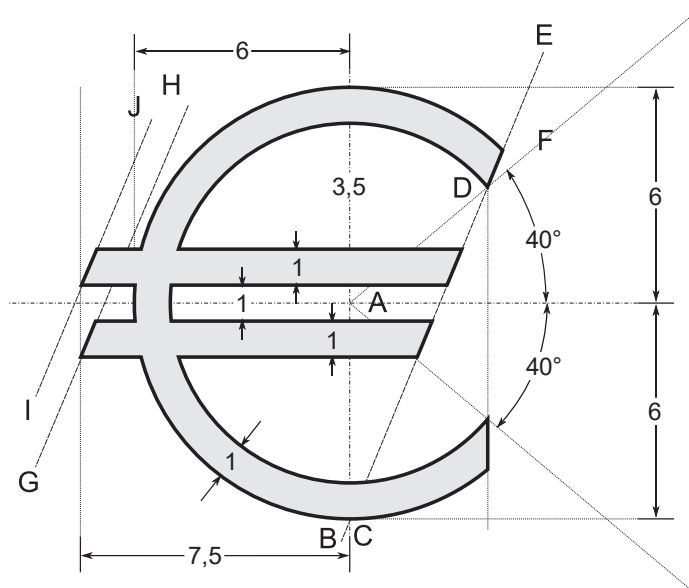
Berendsen besteedt in zijn boek aandacht aan die schrijfwijze, in de Nederlandse als in de Engelse uitgave. Terecht merkt hij op dat die punt in het Engels juist een uitzondering is. Naast China, Japan, Israël en Zwitserland zijn er niet veel niet-Engelstalige landen waar ze ook een punt gebruiken. In veruit de meeste Europese talen en landen is het de komma die als decimaalscheiding dient.

Soms heeft men het over ‘de puntjes op de i zetten’, maar dat is een wat vreemde uitdrukking. Op een i staat namelijk maar één puntje, tenzij het een ï is. En zelfs dat ene puntje staat er al.. tenzij het een 1 is. Maar een 1 is geen i, maar een *puntloze i*, uit het Turkse alfabet. Die komt – net als de i mét punt – voor in de plaatsnaam *Diyarbakır*.

Een punt op een i... Is dat eigenlijk wel een punt? In de afbeelding hieronder zie je de letter i in *Calibri*, *Arial*, *Verdana*, *Times New Roman*, en *Linux Libertine*. Uiterst rechts is geen i, maar een 1 met een punt erop. In *Times New Roman* is een punt een rondje. In *Arial* een vierkant... Nee, wacht: een rechthoek! In *Verdan* ook, maar dan is de rechthoek breder en platter. Sommige mensen vinden het leuk om hierover na te denken. Anderen zeggen: wat is nu eigenlijk je punt?



Tijdens het schrijven van dit hoofdstuk ontdekte ik dat er een officiële norm is voor hoe een euroteken eruitziet. Hij staat in *Wikipedia* en is hieronder afgedrukt. Rechts staat het euroteken in het lettertype van *Meten & Weten*. Eerlijk is eerlijk: dat ziet er net even anders uit. Minder rond en met van die schreefjes. Maar je herkent het wel.



Weten

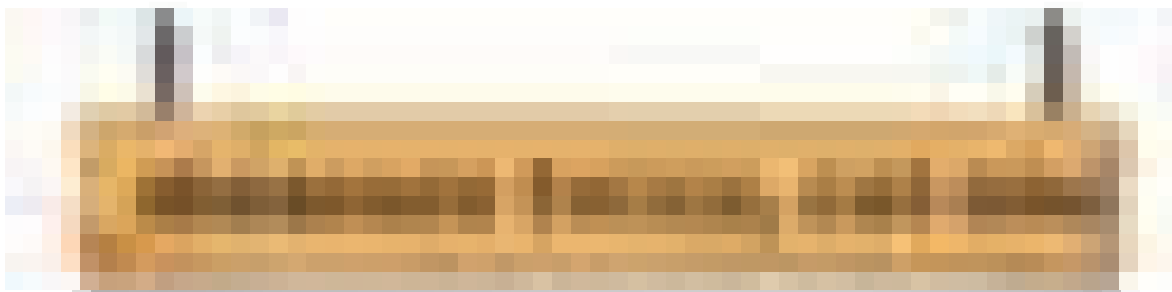
105. Herhalen van een meting

Er is een Engels spreekwoord dat zegt: *Measure twice, cut once*.

Je zou het in het Nederlands kunnen vertalen met: ‘twee keer meten, één keer zagen’. Of ‘snijden’.

Als je de spreuk bij Google intikt, vind je allerlei enorm grappige plaatjes en foto's van bordjes en fotolijstjes met deze tekst, die dan net niet past.

Omdat ik geen foto kon vinden die geinig was en vrij van rechten, en al helemaal geen vectorplaatje, heb ik er zelf eentje gemaakt met een bordje van Pixabay waar je zelf de tekst op kon typen. In een bijpassend lettertype. Daar komt-ie!



De e valt er net af... hilarisch!

Als je bijgekomen bent van het lachen: er zit natuurlijk een serieuze gedachte achter. Wie iets moet opmeten en dan afzagen, is in de aap gelogeed als hij of zij daarna pas erachter komt dat de maat niet klopte. Dan moet je nóg een keer zagen en dat kost je extra tijd en hout. En dus geld. Een tweede keer meten kost ook tijd. Maar meestal minder.

Deze benadering heeft een keerzijde. Als je héél vaak twee keer meet zonder dat je daardoor een fout vindt, kost dat ook veel tijd.

Een belangrijke meting herhalen is verstandig. Maar als je datgene wat je deed op precies dezelfde keer nogmaals doet, is er een reële kans dat je dezelfde fout maakt.

Hoe kun je iets over doen op een andere manier?

- Je zou een andere meetmethode of meetinstrument kunnen kiezen.
- Je zou het iemand anders kunnen laten doen.
- Als beide opties niet mogelijk zijn, zou je kunnen overwegen om het op een andere dag of plaats te herhalen.

De kans dat je iets minstens één keer fout doet, wordt overigens *groter* als je het twee keer doet. Statistisch gezien is de kans dat een handeling die in 90% van de gevallen goed gaat twee keer achter elkaar goed gaat 81%. De kans dat het tien keer achter elkaar goed gaat, is $0,9^{10} \approx 35\%$ Maar dat is niet het sommetje dat we maken. Als je voor 90% zeker weet dat het goed gaat, hoe groot is dan de kans dat het fout gaat? Inderdaad: 10%. En hoe groot is de kans dat het twee keer achter elkaar fout gaat? Inderdaad: 1%.

De enige manier om zeker te weten dat je iets niet fout doet, is: er niet eens aan beginnen. Nul pogingen kunnen niet één keer falen.

Er zijn meer redenen om een meting te herhalen. Eentje die we in dit boek vaak tegenkomen, is dat je door te herhalen een indruk krijgt van de spreiding.

Wat moet je doen als je tien getallen onder elkaar schrijft en de som ervan wilt bepalen? Optellen, inderdaad. En wat als je zeker wilt weten dat er geen fout in de optelling zit?

- Nog een keer optellen.
- Nog een keer optellen, maar dan van onder naar boven.
- Nog een keer optellen, maar dan met een rekenmachine.
- De getallen overtikken in *Excel* en dan de som berekenen.
- Iemand anders vragen om het op te tellen.
- Nog iemand anders vragen om het op te tellen.

Helemaal zeker weten dat je geen fout in een meting of berekening maakt doe je eigenlijk nooit. Maar goede controles helpen wel.

106. Falsificatie

Falsifieerbaarheid of falsificeerbaarheid is een eigenschap van een wetenschappelijke of andere theorie, en wel dat er criteria kunnen worden aangegeven op grond waarvan de theorie zou moeten worden verworpen (niet te verwarren met een theorie die daadwerkelijk is verworpen).

Het falsificationisme is een wetenschapstheorie bedacht door Karl Popper, waarin falsificatie centraal staat.

Bron: Wikipedia

Wat heeft falsificatie met Meten & Weten te maken? In ieder geval heeft het te maken met de gedachte dat je niet zo makkelijk kunt vaststellen dat je iets weet, of beter gezegd: dat je *weet dat iets waar is*. In de wetenschap, software-ontwikkeling en het schrijven van teksten heb je zelden de mogelijkheid om te bewijzen dat iets klopt, laat staan voor 100% foutloos is.

Popper stelde dat je in de wetenschap de dingen zo moet formuleren dat het in principe mogelijk is om de uitspraak te weerleggen als die niet waar is. Een klassiek voorbeeld is: “Alle zwanen zijn wit.” Dat is te weerleggen door één zwaan te vinden die niet wit is. Ga naar het park, zoek naar zwanen en kijk welke kleur ze hebben. Als je een witte vindt, bewijst dat niets. Als je twintig of honderd witte vindt, ook niet. Naarmate je er meer vindt, wordt het *aannemelijker* dat de stelling klopt, maar *bewezen* is het niet. Zodra je één zwaan vindt met een andere kleur, is het klaar. Dan klopt de stelling niet.

Een alternatief is: ga alle voorwerpen en levende wezens af die niet wit zijn. Zodra je er een zwaan tussen treft, is ook bewezen dat de stelling niet waar is.

De stelling klopt ook niet trouwens. Er bestaan ook zwarte zwanen.

Nu de stelling: “Er bestaan ook paarse zwanen.”

Is die te bewijzen? Ja, als je één paarse zwaan vindt, is het nog een twijfelgeval, want de stelling formuleert een meervoud. Maar bij twee is het al bingo: dat is voldoende voor een bewijs.

Is die laatste stelling falsificeerbaar? Nee. Als je je suf zoekt naar zwanen en je ziet nergens een paarse, ga je er misschien wel aan twijfelen, maar bewezen is er niets. Het vervelende is: al zou je je leven lang elke dag overal naar zwanen speuren en je treft nooit een paarse aan, dan heb je nog steeds niet bewezen dat er geen zijn. Misschien bevinden ze zich altijd in een ander land dan jij. Dan kom je ze nooit tegen, al trok je de hele wereld rond.

Het verschil tussen de eerste en tweede stelling zit niet zozeer in de kleur, maar in “alle” en “eentje”. De stelling dat alle zwanen wit zijn, kun je weerleggen door één tegenvoorbeeld. De stelling dat er minstens één bestaat die paars is, kun je alleen maar weerleggen als je *alle* zwanen op de hele wereld gezien hebt. Dat lukt je niet.

Wetenschappers formuleren stellingen die *falsificeerbaar* zijn. Vervolgens doen zij hun best om te *weerleggen* dat ze zelf gelijk hebben. Dat gaat een beetje tegen je gevoel in, want meestal willen mensen gelijk hebben. Maar de kans dat je gelijk hebt wordt juist groter als je je uiterste best doet (samen met anderen) om je ongelijk te halen en daar niet in te slagen.

Softwareontwikkelaars doen dat (als het goed is) ook. Ze schrijven code en in plaats van heel hard “het werkt!” te roepen (dat doen de verkopers wel), doen ze hun uiterste best om een fout te vinden. Gewoon door creatief te zijn in het bedenken van gevallen die mis zouden kunnen gaan. Een tester doet niet zijn of haar best om aan te tonen dat het programma geen fouten bevat. Hij of zij zet zich (binnen zekere grenzen van tijd en middelen) zo goed mogelijk in om juist wél fouten te vinden. En als dat dan echt niet lukt, is het programma misschien wel goed.

Tektschrijvers doen het (als het goed is) ook. Ze schrijven een tekst en in plaats van heel hard “hij is foutloos!” te roepen, speuren ze de regels en alinea’s af om te kijken of er niet een woordje of lettertje ontbreekt. Pas als dat na lang zwijgen, zweten en zwoegen niet lukt, is er een reële kans dat de tekst vrij goed is.

Het lastige van metingen en onzekerheden is dat we nogal eens uitgaan van een normale verdeling. Met een lengte van 0,80 meter en een onzekerheid van 2 cm bedoelen we vaak dat de kans groot is dat de echte lengte tussen 78 en 82 cm ligt, en dat de kans heel groot is dat die tussen 76 en 84 cm ligt. Een grens met een absolute zekerheid is er wiskundig gezien niet.

107. Inschatten wat verwaarloosbaar is

Bij Proeflokaal vandeStreek in Utrecht hebben ze de bierprijs ook naar boven moeten bijstellen. De kosten voor mout zijn flink omhoog gegaan en ook het papier voor etiketten is duurder geworden, zegt commercieel directeur Ronald van de Streek. “In de afgelopen twaalf maanden zijn de kosten zo’n 50 procent gestegen. Toch is het ons tot nu toe gelukt om de prijs van bier slechts met 5 tot 6 procent te verhogen, terwijl de grote merken wel 15 procent duurder werden. Daarmee zijn we in prijs dichterbij de grote brouwerijen in de buurt gekomen.”

Bron: Trouw, 27 januari 2023

De prijs van een fles bier hangt van allerlei factoren af. De kosten van mout en gerst bijvoorbeeld, maar ook van water. En van glas en papier, en metaal voor de kroonkurk. Arbeidsloon en energiekosten zijn ook al van invloed, net als accijnzen, reclamespots en transportkosten.

Bij veel berekeningen verwaarloos je bepaalde invloeden. Dat is handig, omdat het alles wat *eenvoudiger* maakt. Verwaarlozing zorgt ervoor dat je snel kunt *schatten*. En hoewel schatten veel minder zekerheid verleent dan berekenen, geeft het door de eenvoud en snelheid wel veel inzicht.

Om te weten hoe de kosten van een flesje bier afhangen van de stijging van de papierprijs, moet je een aantal dingen weten.

1. De prijs van papier.
 - a. Hoeveel was het eerst?
 - b. Hoeveel is het nu?
2. Hoeveel papier er op één fles zit.
 - a. De oppervlakte van de etiketten bij elkaar.
 - b. Het gewicht (de massa) per oppervlakte.

‘Ergens op het internet’ staat dat papier in prijs stijgt van € 50 naar € 75 per ton en dat de afmetingen van een etiket 180 mm × 90 mm zijn. Het etiket op de achterkant is meestal kleiner en vaak zit er nog wel wat papier om de hals. Een stijging € 50 per ton en een oppervlakte van 500 cm² is dus een ongunstige schatting. Het papier zal minder wegen dan 100 g/m².

Laten we werken met wetenschappelijke notatie: eerst alles omzetten euro’s en SI-eenheden, machten van 10 en één significant cijfer per grootte.

Prijsstijging:	€ 50 per ton	=	$5 \cdot 10^{-2}$ euro/kg.
Oppervlakte van het papier:	0,05 m ²	=	$5 \cdot 10^{-2}$ m ²
Massa papier per oppervlakte:	0,1 kg/m ²	=	$1 \cdot 10^{-1}$ kg/m ²

Het lijkt overdreven om deze stappen zo lang gescheiden te houden en zo strikt in eenheden en machten van 10 te noteren, maar naarmate je daar voorzichtiger in bent, is de kans op fouten kleiner. Een nulletje meer of minder, de komma verkeerd zetten: op een tentamen zijn het tienden in je cijfer. In de werkelijkheid van de ingenieur kan het tonnen aan schade betekenen. Getallen met vijf of meer cijfers (en een euroteken) voor de komma.

Nu kunnen we deze drie getallen met elkaar vermenigvuldigen. We komen dan uit op $25 \cdot 10^{-5} = 2,5 \cdot 10^{-4}$ euro.

In een artikel, verslag, e-mail of app-bericht is het handig om nog één slag te maken, namelijk je afvragen hoe het voor de doelgroep het meest overtuigend overkomt. Die $2,5 \cdot 10^{-4}$ euro is voor veel lezers niet handig. Je kunt die noteren als “minder dan 0,1 eurocent” of als € 0,001.

Hoe kwamen we hier ook alweer op? Als je een antwoord hebt, moet je je altijd afvragen wat de eigenlijke vraag ook alweer was. Dat was *niet* wat een los etiket kost; het ging om de invloed van de prijsstijging van het papier op een fles bier. En de concrete vraag was of het *verwaarloosbaar* is of niet.

Het antwoord: ja, dat is verwaarloosbaar. Tenzij je héél goedkoop bier hebt.

Oefening

1. Stelling: “Dit boek bevat 366 pagina’s.”
 - a. Is deze stelling juist?
 - b. Wat denk je dat de auteur gedaan heeft om het (al dan niet) juiste aantal in de stelling op te nemen?
 - c. Als er op elke drie pagina’s van dit boek één spelfout staat, hoeveel procent van de woorden bevat er dan een spelfout?
2. Dit boek bevat 366 pagina’s en is tegen kostprijs te koop. Dat bedrag is te berekenen via de website van Pumbo.nl. Maak een schatting van de bijdrage van de papierprijs in de kostprijs van dit boek en geef aan of dit verwaarloosbaar is.

108. Vroeg vinden van verbeter(d)ingen

Software-ontwikkelaars zijn bekend met de wetmatigheid dat een fout duurder wordt naarmate je die later vindt. Je tikt een regel code in, kijk ernaar, en ziet dat er iets niet correct is. De schade valt dan mee. Een keer met je muis klikken en de *Backspace* indrukken en je bent er al bijna. Dat kost een minuut van je tijd en dus weinig geld.

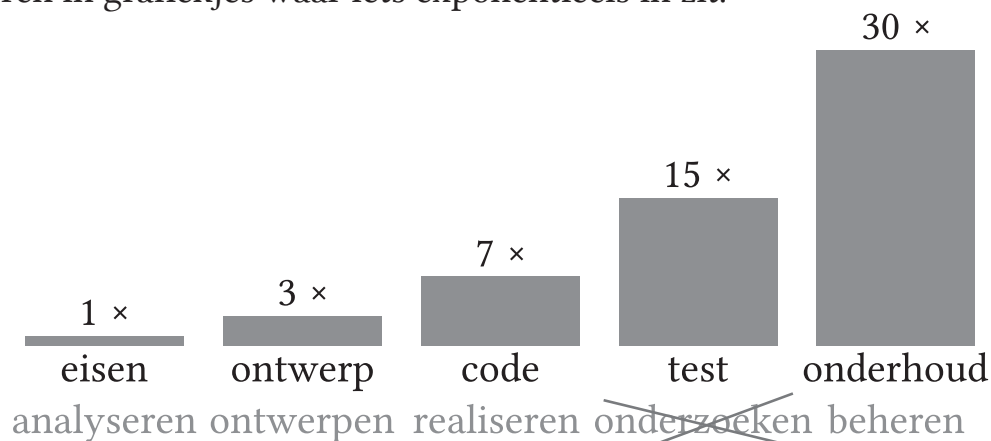
Als je de fout vindt tijdens de controles die je doet voordat je de software oplevert aan een tester, is het al wat duurder. Je moet dan langer zoeken naar waar het zit. Je bent dan gelukkig nog steeds de enige die er tijd aan kwijt is.

Als de tester 'm vindt, gaat er meer tijd in zitten. Die heeft er zelf werk aan en moet ook nog eens een *bug report* schrijven. Daar moet de programmeur dan ook weer wat aan toevoegen en bij de volgende oplevering aan de tester moet bevestigd worden dat de bug eruit is.

Als een tester een bug vindt, moet je nooit boos op hem of haar worden. Een dikke zoen is – als de persoon in kwestie dat op prijs stelt! – een betere reactie. Een bosje bloemen, een kop koffie, een doos *Merçi*: wees blij dat de tester die fout gevonden heeft.

Natuurlijk kost het je tijd om dan die fout te gaan oplossen. Maar het is niet zo dat de fout er niet in zou zitten als hij niet gevonden was. De honderden of duizenden gebruikers van je software zijn veel meer uren met dat programma bezig. Als die de fout vinden, zijn ze soms niet blij. Omdat ze er last van hebben. Ze kunnen een schadeclaim indienen. Of naar de concurrent overstappen. In het gunstigste geval melden ze het beleefd. En dan ben je veel tijd kwijt om het alsnog op te lossen..

Er zijn wel onderzoeken naar die toename van de kosten gedaan. Die resulteren in grafiekjes waar iets exponentieels in zit.



Dit plaatje is gebaseerd op gegevens van het *National Institute of Standards and Technology*. De rechter kolom zou eigenlijk hoger moeten zijn. Het NIST stelt namelijk dat die laatste een factor 30 tot 100 groter is dan de eerste.

Onder de bekende fases van software-ontwikkeling heb ik de competenties van een mechatronics gezet. Het klopt niet één-op-één en voor testen is er niet een aparte competentie. Eigenlijk hoort testen gewoon bij het realiseren (schrijven) van code.

Geldt wat hier over software staat ook voor metingen? Zijn de kosten van een meetfout ook hoger als je die niet meteen zelf ontdekt? Ja. Het is alleen wat lastiger in kaart te brengen.

- Wat als je een fout maakt in een meting of een berekening of een bepaling van een onzekerheid en het wordt niet ontdekt voordat je afstudeerverslag ingeleverd wordt, maar tijdens de examenzitting?
- Wat gebeurt er als een metalen constructie veel te sterk wordt uitgevoerd als gevolg van een verkeerde berekening?
- Wat zijn de gevolgen als een patiënt een factor duizend te veel morfine krijgt toegediend? (“Echt, alleen de komma stond verkeerd! Op school kreeg je dan een paar tienden puntenaftrek...”)

Er zit ook een andere kant aan het verhaal. Ja, een fout is inderdaad goedkoper op te lossen naarmate je die eerder vindt. Maar er is ook een soort omgekeerde wet.

Door heel even een stuk programmacode, een tekst of een berekening na te kijken, vind je al heel snel een fout.

- Door langer te kijken vind je meer fouten, maar niet heel veel meer.
- Door heel lang te kijken vind je er nog meer.
- Om de laatste fout te vinden, moet je extreem lang zoeken.

Hoe eerder je de fout vindt, hoe goedkoper het is. Maar als je foutloos wilt werken, gaat je werk wel heel erg duur worden.

Een student zei eens: liever een zes en geen stress, dan een zeven en geen leven. Dat rijmt. Of de afweging terecht is, is een andere vraag. Studenten die altijd maar op een 6 mikken, komen in hun afstudeerwerk minder beslagen ten ijs dan anderen die probeerden om minstens een 8 te halen. Die gaan met minder stress hun examenzitting in. Hebben gewoon meer bagage dan anderen.

Van wie gaat voor een acht, kan veel goeds worden verwacht.

109. Proberen

“Het woord ‘proberen’ telt zeven letters.”

Wat doe je met deze uitspraak? Voor waar aannemen? Of ga je hem toetsen?

Als je hem voor waar aanneemt: gefeliciteerd! Je bent heel snel klaar met dit hoofdstuk. Voor wie de uitspraak wil toetsen: het tellen van de letters is de aangewezen oplossing. P-R-O-B-E-R-E-N. Je kunt dat op de vingers van twee handen doen. Of typt het woord over in een tekstverwerker (kopiëren uit een digitale versie van de tekst is nog beter) en laat die het werk doen. Met weinig moeite vind je via Google een website die het voor je doet. “Maar ik kan toch heus wel tot tien tellen?” Ja. Maar als je iets *helemaal zeker* wilt weten, volsta je niet met één keer meten. Iets wat echt belangrijk is, stel je op minstens twee manieren vast. Het liefst door twee verschillende personen.

Er is nog iets. We hebben niet eenduidig vastgelegd wat we onder ‘bevatten’ verstaan. Als één letter twee keer voorkomt, tellen we die dan dubbel? Een definitieprobleem komt ook voor onder zeer eenvoudige vragen. Misschien wel vaker dan onder moeilijke vragen.

Voor wie het nog niet nageteld had: proberen bevat acht letters. Maar dan tel je de E dubbel, dus volgens een bepaalde definitie zijn het er toch zeven. Mis! De R komt namelijk ook twee keer voor, maar die valt minder op.

“Het woord ‘proberen’ telt vijf verschillende letters.” Is dat waar? Ja. Dat kun je op je vingers natellen: P-R-O-B-E-R. De e die erna komt, verschilt van die eerste e en de N op het eind hadden we ook nog niet gehad. Maar de vraag was niet of het woord *precies* vijf verschillende letters bevat. Als het er zes of zeven zijn, zijn het er ook vijf. Meer nog zelfs.

“Goedemorgen, bakker.”

“Goedemorgen, groetende klant die binnenkomt.”

“Hebt u misschien vijf broden?”

“Nee, helaas. Ik heb er nog zes. Of nee: zeven.”

“Helpt u dan die mevrouw na mij eerst maar. Die wil er namelijk twee.”

“Dat zou mooi zijn, want dan heb ik er nog vijf.”

“Hebt u er dan nóg vijf? Dat zouden er dan tien zijn...”

(De mevrouw achter deze klant verlaat nu hoofdschuddend de winkel.)

Het volgende belangrijke aspect: omdat die ‘zeven’ al genoemd is, blijft dat getal in je hoofd hangen. Als iemand dan ook nog eens begint over *één* letter die dubbel voorkomt, dan zou je in de verleiding kunnen komen om te denken: ‘nu snap ik waar ze zeven zeiden!’

Het laatste aspect noem je ook wel anchoring. In het Nederlands *verankerings*effect of *referentie-effect*. Het ankerpunt is datgene waarvan je uitgaat en waaraan je in zekere zin vastzit.

“Zijn er meer of minder dan tweehonderd horecazaken in Delft?”

Je denkt na over de vraag, maakt wat afwegingen in je hoofd en kiest uit één van beide mogelijkheden: meer of minder. (Het kunnen er ook exact tweehonderd zijn, dan zou je de vraag alleen maar goed hebben als je ‘nee’ antwoordt. ‘Ja’ is trouwens wel goed, want er zijn inderdaad niet precies tweehonderd horecazaken in Delft. Het zijn er meer of minder.)

Nu de volgende vraag: “Hoeveel denk je dat het er zijn?”

Uit wetenschappelijk onderzoek blijkt dat schattingen gemiddeld genomen *lager* uitpakken dan wanneer de eerste vraag het getal driehonderd bevat zou hebben. Er is dus een (waarschijnlijk onbewuste) neiging om dat getal als uitgangspunt te nemen. “Het zal toch niet voor niets genoemd zijn...”

Voor metingen geldt ook zo’n soort effect. Naast *anchoring* kun je trouwens ook last hebben van het feit dat je het echte antwoord al weet. Wie de versnelling van de zwaartekracht wil meten en op $3,2 \text{ m/s}^2$ uitkomt, denkt meestal dat hij of zij iets fout gedaan heeft. Dan doe je de meting en/of de berekening over totdat je in de buurt komt van $9,81 \text{ m/s}^2$. Tenzij je extreem veel zelfvertrouwen hebt en denkt dat de hele wereld ongelijk heeft.

“Het woord ‘proberen’ telt zeven letters.” Kunnen woorden dan tellen? Dat *woord* telt de letters toch niet? Die bakker telde (en vertelde) hoeveel broden er waren. Volgens *Van Dale* heeft ‘tellen’ vijf verschillende betekenissen. Eentje daarvan is ‘bevatten’. Soms is taal niet te bevatten.

110. Pas op

Het schijnt zo te zijn dat veel lezers over het hoofd zien dat ‘pas op voor de hond’ een fout bevat.

Antwoorden

Hoofdstuk 0

1. Elke meting heeft een onzekerheid. Daar moet je een schatting van geven, omdat de uitkomst van de meting anders weinig betekenis heeft. Nadenken over hoe je die onzekerheid kleiner kunt krijgen, kan nuttig of zelfs nodig zijn.
2. Niet alles in één keer lezen. Oefeningen doen. Erover praten met anderen, de discussie in de lessen volgen.

Hoofdstuk 1

1. Nee. Als de meter in 1986 een centimeter *langer* was dan tegenwoordig, zou dat wél een verklaring kunnen zijn waarom jongeren nu zo'n anderhalve centimeter langer zijn volgens "de nieuwe meetlat". Als een meter vroeger 'meer' was, was je zelf volgens die meetlat korter.
2. In 1983.
3. Met het schaamrood op de kaken: van *Wikipedia*.
Dat zal heus wel kloppen, maar kun je in je afstudeerverslag niet maken en dus als docent in een les of boek ook niet doen.
Op de *Wikipedia*-pagina zelf kun je wel betrouwbare bronnen vinden.
4. Nee. De schrijver bedoelt een afwijking van 2σ in beide richtingen.
Je mag maximaal 2σ "te klein" zijn en maximaal 2σ "te groot".
5. Twee aspecten:
 - a. De waarde 1,414 213 562 – het zijn trouwens de eerste decimalen van $\sqrt{2}$ – heeft veel meer cijfers dan redelijk is. "Anderhalve seconde" zou zinniger zijn.
 - b. Dat je de galm wel of niet meerekent, moet je vermelden om een definitieprobleem te voorkomen.
6. Het spreekt niet vanzelf, een eenheid moet je altijd vermelden. Ook als die voor de hand ligt.

Hoofdstuk 2

1. Verbeteringen:
 - a. "Om met 100% zekerheid te kunnen garanderen dat iemand te hard heeft gereden" moet zijn: "Om te kunnen garanderen dat iemand *alleen wordt bekeurd als* hij te hard heeft gereden"
 - b. "trekt de politie van de gemeten snelheid een paar kilometer af" moet zijn "trekt de politie van de gemeten snelheid een paar km/u af"
 - c. "3 kilometer meetcorrectie" moet zijn: "3 km/u meetcorrectie"
 - d. "Er geldt een ondergrens van 4 km/u voordat wordt bekeurd" moet zijn: "Er wordt pas bekeurd als er minstens 4 km/u te hard wordt gereden."
 - e. Het is niet duidelijk of een snelheid precies op de ondergrens als te snel wordt beschouwd.

- f. Strikt genomen wordt er niets gezegd over een snelheid van exact 100 km/u, alleen over meer of minder.

Logisch gezien is dat geen probleem, want 3% is bij die snelheid precies gelijk aan 3 km/u.

2. Dit is een gemene vraag, maar er staan in *Metten & Weten* wel vaker gemene vragen, omdat het om de *oefening* gaat en niet om het goed hebben van het antwoord. Het is geen tentamen: daar horen geen gemene vragen te worden gesteld.

Wat is er gemeen aan?

- a. Er staat niet bij wat de maximumsnelheid is, dus je kunt de vraag helemaal niet beantwoorden. Sommige lezers zullen automatisch uitgegaan zijn van 100 km/u als maximumsnelheid, maar het is niet gegeven.
- b. Er wordt gesproken over *iemand die 107,2 km per uur rijdt*. Als dat de *echte* snelheid is, hoeft er niet gecorrigeerd te worden. Als dat de *gemeten* snelheid is, moet er eerst 3% worden afgetrokken. Maar er staat niets over 'gemeten'.
Mocht je de vraag toch zo opvatten: na correctie van 3% kom je op 103,984 km/u. Dat is (bij een maximumsnelheid van 100 km/u) minder dan 4 km/u te snel.

Hoofdstuk 3

1. Alle sleutels zijn even waarschijnlijk. De kans dat de eerste goed is, is $1/8$. De kans dat een van de eerste twee goed is, is $2/8$. Een van de eerste drie: $3/8$. Een van de eerste vijf: $5/8$. (Een van alle acht: $8/8$).
2. De redeneerfout is dat weliswaar alle mogelijkheden worden genoemd (wel of niet de juiste) maar dat de kansen over beide mogelijkheden niet gelijk verdeeld zijn. (Volgens deze redenering is de kans op de hoofdprijs in de Staatsloterij ook 50%. Wat niet zo is.)
3. Het moet toch gek lopen willen alle 389 studenten op een *andere* dag jarig zijn... Er zijn maar 366 verschillende verjaardagen mogelijk. We kunnen 367 of meer studenten meer zo verdelen dat iedereen op een andere dag jarig is. Zoveel dagen zijn er namelijk niet. (De kans dat iemand op dezelfde dag *als jij* jarig is, is overigens wel kleiner dan 100%. Maar dat was de vraag niet.)

Hoofdstuk 4

Belangrijk is de vraag of je een van beide partijen bevoordeelt. Wat is erger? Dat het een doelpunt was en toch niet telt, of dat het geen doelpunt was terwijl hij wel telt?

Hoe dan ook gelden de regels / het systeem voor beide ploegen.

Hoofdstuk 5

1. Alleen die gedachte is fout. Het goede nieuws: zodra je dat weet, is de fout ook weg. Het probleem dat sommige mensen hebben, is dat ze zichzelf *identificeren* met hun

gedachten, meningen en standpunten. Maar je *bent* niet je gedachten, meningen en standpunten, je hebt ze. Als jouw gedachte (of mening of standpunt) fout is en je weet het, dan ruil je die gewoon gratis in voor een betere. Je raakt jezelf niet kwijt, maar wel iets waar je liever van af bent.

2. Dit is helemaal geen dyslexie-test. Het is gewoon een flauwe vraag. Maar wie het weten wil: het verschil tussen de woorden *meting* en *meting* zit in de letter i van het tweede woord. Dat is niet 12-punts, maar 11,5-punts.
3. Met de tekens *liggend streepje*, *schuin streepje*, *haakje sluiten* en *punt* moet je het al kunnen redden.
4. Ik weet niet waarom je deze oefeningen gelezen (en misschien wel gedaan) hebt, terwijl er boven de opdracht stond dat ze geen enkel nut hebben. Misschien kun je het zelf bedenken.

Als je toch bezig bent met nutteloze dingen: dat schouderophalen-plaatje staat ook ergens verstopt in de geodriehoek op de omslag van dit boek. Heel klein ergens.

Hoofdstuk 8

1. Box zei meer dan dat alle modellen fout zijn: hij zei ook dat sommige nuttig zijn. We gebruiken de modellen die nuttig zijn.
2. Het gaat om de eenzijdige overschrijdingskans bij vijf keer de standaardafwijking. Die bereken je eenvoudig met *Excel*, met de formule

`=NORM.VERD(0;1,5;0,3;WAAR).`

De kans is ongeveer $3 \cdot 10^{-7}$.

Erg klein dus, maar als alle mensen op deze wereld één zo'n meting doen, zou dit resultaat best kunnen voorkomen.

Hoofdstuk 9

1. $2,7 \pm 0,3$ V
2. $2,7(3)$ V

Hoofdstuk 11

1. Thermometers 6 en 7 laat je buiten beschouwing, want die zijn niet goed. Van de andere vijf zou je in het beste geval het gemiddelde kunnen nemen. Dat verkleint de toevallige fout. Eenvoudiger (en praktisch even goed) is de middelste thermometer nemen. Die heeft ongeveer het gemiddelde vloeistofniveau.
2. Deze vraag kun je alleen maar beantwoorden als je meer weet over de tolerantie die gekozen is. Als de fabrikant weet dat er 5% spreiding in kan zitten en dat acceptabel vindt én in de specificaties vermeldt, dan is er sprake van een onzekerheid. Als de norm is dat er maar 2% afwijking mag zijn, dan valt deze buiten de norm. Dan hoort het product niet verkocht te worden.

Hoofdstuk 12

1. Zie figuur op Nederlandstalige pagina van *Wikipedia*.

2. Zie figuur op Engelstalige pagina van *Wikipedia*.
3. Eerst het flauwe antwoord: ja. Je kunt op *elke vraag* een exact antwoord geven.
 Vraag: "Hoeveel inwoners heeft China?"
 Antwoord: € 16,42. Dat is exact, maar het slaat nergens op...
 Nu het normale antwoord. Gevraagd wordt om een afronding op gehele inches. Dat is 11 802 852 677, een getal op met elf cijfers. Weliswaar afgerond, *maar dat was de vraag ook*. Deze elf cijfers zijn exact.
4. Het kan natuurlijk niet helemaal. Op twee na kun je de letters een plekje geven. Je zou van wat je over hebt (de d en de g) een c en een o kunnen maken met wat kras-sen. (Misschien een flauwe oefening, maar nu vergeet je het tenminste nooit meer.)

Hoofdstuk 13

1. Strikvraag! Er wordt gevraagd naar een *nauwkeurigheid*, maar wat geven is, is de *onzekerheid* van twee multimeters. Die onzekerheid bestaat uit een toevallig deel en een systematisch deel en je weet niet welke van beide het meeste bijdraagt. Dus kun je ook niet veel zeggen over de nauwkeurigheid. (Hooguit dat die niet groter kan zijn dan een bepaalde waarde. Voor de tweede is de onzekerheid kleiner, dus als je moet gokken, neem je die. Maar je weet het niet.)
2. Strikvraag! Er wordt alleen iets gezegd over de spreiding. Die uitspraken hebben dus betrekking op de precisie, niet op de nauwkeurigheid.

Hoofdstuk 15

1. We hebben eigenlijk veel te weinig informatie, maar laten we uitgaan van een standaard meetlint / rolmaat. Die heeft een millimeteraanduiding. Theoretisch zou je dan op halve millimeters kunnen aflezen. Dat doe je aan twee kanten (bij de nul en bij de waarde die je afleest) en de afleesfout zou dus 1 mm zijn. Lukt dat over een afstand van ruim 2,5 m? Je hebt misschien ook nog te maken met een ongelijke ondergrond. Daarom is 2 mm of zelfs 5 mm reëler.
 In praktijk zou het aardig zijn om verschillende mensen dezelfde meting te laten doen. De spreiding geeft dan een aardig beeld van de werkelijke afleesfout.
2. De steekproef-standaardafwijking is afgerond op drie significante cijfers gelijk aan 0,0626 m. Het is gebruikelijk om die als aanduiding voor de spreiding in de meting zelf te geven.
Je ziet in deze opgave (en in dit specifieke voorbeeld) dat de spreiding in datgene wat we meten een factor tien groter kan zijn dan de onzekerheid ten gevolge van de meting zelf.

Hoofdstuk 16

1. Het schatten van de hoek is niet eens zo moeilijk. Gewoon een geodriehoek erlangs leggen is een optie. Beter is om de hoogte en de breedte van de driehoek te meten. Dat is makkelijker aflezen, maar je moet dan nog wel een rekenmachine pakken voor die hoek.
 Als je gebruik maakt van een PDF (niet een gedrukt boek of een print) dan zou je

ook op *Print Screen* kunnen drukken en de schermafdruck zodanig uitknippen dat je de hoogte en breedte in pixels kunt opvragen. Mijn eigen aanpak leidde tot 615×188 pixels.

Lastiger is de vraag wat de onzekerheid is. Dat is de essentie van deze oefening. Iets zo goed mogelijk doen is makkelijker dan weten hoe goed je het doet. Zelf denk ik dat drie pixels in hoogte en breedte redelijk is. Iemand anders zal één pixel zeggen, of vijf...

De twee uiterste hoeken krijg je dus voor

- $612 \times 191 \text{ pixels} \rightarrow \tan \theta = 191/612 \rightarrow \theta \approx 17,3327^\circ$
- $618 \times 185 \text{ pixels} \rightarrow \tan \theta = 185/618 \rightarrow \theta \approx 16,6652^\circ$

De helft van het verschil tussen deze waarden zou een redelijke schatting zijn van de onzekerheid. Afgerond op één significant cijfer kom je dan op $0,3^\circ$.

2. Deze vraag is eigenlijk niet te beantwoorden. Als onze reactietijd altijd exact 0,3 seconde zou zijn, dan ben je zowel bij het starten als bij het stoppen van de tijdsmeting exact 0,3 seconde te laat. Het tijdsinterval zou dan geen onzekerheid hebben: die twee keer te laat drukken heft elkaar precies op.

Eigenlijk zou je moeten weten hoe groot de *spreiding* is in onze reactietijd. Die kun je bepalen door de meting vaak te herhalen.

Hoofdstuk 25

1. Dit is een wat vreemde opgave, maar het gaat om de *oefening*. De bedoeling is dat je heel goed nadenkt over de *exacte* betekenis van de begrippen.

- a. *measurement* (meting)

Deze lijkt nogal makkelijk: elke meting heeft een onzekerheid.

Maar een meting bestaat uit een beste schatting én de onzekerheid. De beste schatting zelf is niet de meting.

Voorbeeld: we meten een tijd en komen op $2,6 \pm 0,3 \text{ s}$.

Hierin is 2,6 s de *beste schatting* en 0,3 s de *onzekerheid*.

De meting is de combinatie van beide.

Wat is nu het antwoord op de vraag? “Eet je appeltaart met slagroom *met slagroom*?” Ja. Tenzij je bedoelt dat je bij appeltaart met slagroom apart slagroom bestelt. Dan kan wel, maar is niet gebruikelijk. Je hebt dan een dubbele portie.

- b. *true value* (werkelijke waarde)

Deze vraag is wat eenduidiger. De *true value* (als die bestaat) heeft geen onzekerheid. Het is een exact getal, dat je meestal niet kent. Dat onbekende getal zelf heeft geen onzekerheid.

- c. *best estimate* (beste schatting)

Zie onderdeel a hierboven. Je zou kunnen zeggen dat de beste schatting van 2,6 s een onzekerheid heeft van 0,3 s. Maar je kunt ook zeggen dat die beste schatting gewoonweg de uitkomst van een sommetje of meting is.

Wat je afleest van het meetinstrument of de rekenmachine kan in principe volstrekt ondubbelzinnig zijn.

- d. *accuracy* (nauwkeurigheid)
precision (precisie)
uncertainty (onzekerheid)

Inhoudelijk zijn deze drie het meest interessant.

Een onzekerheid bestaat uit twee delen: een deel dat een gevolg is van een systematische afwijking (nauwkeurigheid) en een deel dat een gevolg is van een toevallige afwijking (precisie). Voor alle drie geldt dat het berekende en/of geschatte waarden zijn. Die hebben zelf ook een onzekerheid.

- e. *accepted value* (geaccepteerde waarde)
Belangrijk is om het verschil in te zien tussen *true value* en *accepted value*. De *accepted value* van bijvoorbeeld een natuurconstante wordt nogal eens gebruikt als de ‘echte’ waarde, omdat je de *true value* niet kent.

Hoofdstuk 26

1. Om dat te kunnen uitrekenen, moet je eerst bedenken hoeveel mogelijke cijfers er zijn. Alle cijfers van 1 t/m 9 kunnen tien verschillende decimalen (0 t/m 9) hebben. Alleen een 10 wordt altijd gevolgd door een 0. Er zijn dus $9 \times 10 + 1 = 91$ mogelijke cijfers. De verschillen zitten in alle even cijfers die gevolgd worden door een 5. Dat zijn dus de 2,5, de 4,5, de 6,5 en de 8,5. Vier van de 91 cijfers pakken met ‘gewoon’ afronden precies een punt hoger uit dan met *banker's rounding*. Het gemiddelde voordeel dat studenten hebben met afronding is dus $4/91 = 0,044$ punt. Is dat veel? Het valt wel mee, maar er zijn in Nederland meer dan achthonderdduizend studenten aan hogescholen en universiteiten. Als die allemaal twaalf cijfers per jaar halen, dan worden er jaarlijks toch bijna een half miljoen punten cadeau gegeven. Daar hoor je nou nooit eens iemand over klagen...
Voor studenten die het nog niet helemaal begrijpen en bang zijn dat de afronding ooit zal veranderen: een 5,5 is sowieso een 6.
2. De enige situatie waar het uitmaakt, is het cijfer 5 dat voorafgegaan wordt door een even cijfer. In dat geval gaat *banker's rounding* omlaag (naar dat even cijfer toe dus) en gewoon afronden omhoog.
In het getal 3,14159265358 staan drie vijven. Alleen de tweede wordt voorafgegaan door een even cijfer. Bij afronden op zeven decimalen krijg je dus een verschil: 3,1415926 voor *banker's rounding* en 3,1415927 voor gewoon afronden.
3. Voorwaarde om deze opgave te kunnen maken, is dat je snapt hoe kommagetallen werken als je hexadecimaal rekt. Het gaat niet anders dan bij decimaal rekenen. Alleen staat werk je nu niet met 10^{-1} en 10^{-2} voor het eerste en tweede cijfer achter de komma, maar met 16^{-1} en 16^{-2} .
Ook het afronden gaat op dezelfde manier. Kijk naar het cijfer dat moet “verdwijnen”. Dat is een 8. Een 8 is te vergelijken met de 5 bij decimaal tellen: het is namelijk de helft van het grondtal. Een 8 wordt naar boven afgerond bij gewoon afronden en naar het even cijfer bij *banker's rounding*. Is het cijfer ervoor even of oneven? Een hexadecimale F is oneven. (Het is namelijk niet deelbaar door 2.) Het afronden van 1A,F84 gaat dus in beide gevallen naar boven.

4. Een binair kommagetal werkt óók gewoon als een decimaal getal. De cijfers achter de komma hebben als waarde niet 10^{-1} en 10^{-2} en 10^{-n} , maar 2^{-1} en 2^{-2} en 2^{-n} . Dat getal 101,01 is dus gewoon gelijk aan 5 (die 101 voor de komma) plus 0,25 (die 01 na de komma).

Nu het afronden. Je hebt maar twee cijfers, namelijk 0 en 1. Die 1 is te vergelijken met de 5 bij een decimaal talstelsel. Het is namelijk de helft van het grondtal.

Wat misschien een beetje raar “aanvoelt”, is dat een binair talstelsel maar twee cijfers heeft. Bij een decimaal talstelsel heb je 0 t/m 4 (die omlaag worden afgerond) en 5 t/m 9 (die omhoog worden afgerond). Bij een binair talstelsel bestaat de eerste helft van de cijfers uit alleen die 0 en de tweede helft uit alleen die 1. Bij gewoon afronden wordt elke 0 omlaag afgerond en elke 1 omhoog. Bij *banker's rounding* wordt die 1 afgerond naar het even cijfer... naar de 0 dus.

Het principe is niet anders dan bij decimale getallen. Sterker nog: het is misschien wel duidelijker. Een 0 hoeft je niet af te ronden: die laat je gewoon weg. Een 1 moet je wel afronden en dat wil je dus in de helft van de gevallen omlaag doen en in de andere gevallen omhoog.

Het getal 101,01 wordt in één decimaal dus 101,1 bij gewoon afronden en 101,0 bij *banker's rounding* wordt het 101,0.

Hoofdstuk 27

1. 1
2. 1
3. 1
4. 2
5. 2
6. 3
7. 1 (en dat ‘voelt’ een beetje raar, want je geeft het geldbedrag net zo nauwkeurig weer als in de vorige oefening)
8. 4
9. 1, 2, 3, 4, 5, 6 of 7 (en die niet-eenduidigheid zou te voorkomen zijn door wetenschappelijke notatie)
10. 2 (en meteen een voorbeeld van de wetenschappelijke notatie)
11. 9 (of misschien toch minder?)

De punten in het getal maken niet uit, die zijn alleen voor de leesbaarheid. Maar het feit dat er twee nullen achter de komma staan suggereert dat die andere nullen ook significant waren.

Bij de laatste oefeningen ging het telkens om ‘een miljoen’. De wetenschappelijke notatie is eenduidig.

Hoofdstuk 28

1. 460
2. 0
3. 314
4. 4,72
5. $1 \cdot 10^2$

6. 1,5240
7. De lichtsnelheid is exact gedefinieerd als 299 792 458 m/s.
Een uur is exact $60 \times 60 = 3600$ seconde.
Een vermenigvuldiging levert op: km/u.
Hoeveel significante cijfers zijn er? Negen? Daar lijkt het op, maar zowel voor de getallen 299 792 458 als 3600 geldt dat ze exact zijn.

De exacte uitkomst is 1 079 252 848 800 km/u.

Hoofdstuk 29

1. 6 s
 2. 15 m
 3. $1,6 \cdot 10^6$ kg
 4. 3 A
 5. 1,9 K
 6. 0,2 V
- Door de afronding wordt het eerste cijfer een 2 en dus geen 1.
Daarom noteren we maar één significant cijfer.
Het is verleidelijk om er 0,19 V van te maken, maar dan klopt je afronding niet.

Hoofdstuk 30

1. lengte = $1,710 \pm 0,015$ m
2. breedte = 200 ± 13 mm
3. oppervlakte = $60,55 \pm 0,05$ m²
4. volume = $403,020 \pm 0,010$ ml
5. tijd = $59,988 \pm 0,005$ s
6. massa = $72,6 \pm 0,4$ kg
7. hoek = $3,14 \pm 0,03$ rad
8. volume = $98,4 \pm 0,2$ cm³
9. frequentie = 50 ± 3 Hz
10. temperatuur = 20 ± 4 °C
11. elektrische stroom = $1101,0 \pm 1,5$ mA
12. volume = $(4400 \pm 7) \cdot 10^3$ μm³
13. spanning = $230,00001 \pm 0,00005$ V
14. elektrische stroomdichtheid = $0,0 \pm 0,4$ A/m²
15. warmte = $(6,0 \text{ kJ} \pm 0,4) \cdot 10^3$ kJ
16. oppervlakte = $7659,0 \pm 0,9$ nm²
17. kracht = 100 ± 10 kN
18. vermogen = $0,22 \text{ kW} \pm 0,03 \text{ kW}$
19. druk = $200,00 \pm 0,03$ N/m² (of Pa)
20. weerstand = $20,0 \text{ k}\Omega \pm 0,3 \text{ k}\Omega$

Hoofdstuk 31

De juiste notatie is: $9,84 \pm 0,18 \text{ m/s}^2$.

Er mogen ook haken omheen: $(9,84 \pm 0,18) \text{ m/s}^2$.

De zeven opgaven samen bevatten *dertien* fouten. Had je ze *alle dertien*?

1. verkeerde significantie: fout moet in twee decimalen, want het eerste significante cijfer is een 1
2. onzekerheid verkeerd afgerond
3. onzekerheid heeft een significant cijfer te veel
4. beste schatting is verkeerd afrond én er ontbreekt een spatie voor de eenheid
5. punten gebruikt in plaats van komma's
6. beste schatting heeft een significant cijfer te weinig
7. punten in plaats van komma's
verkeerd afgerond
geen onzekerheid opgegeven
geen eenheid opgegeven
significante cijfer te veel
8. Schrijf eerst de formule zo dat G links van het isgelijktteken komt te staan.

$$G = F \frac{r^2}{m_1 \times m_2}$$

Daarna kun je de eenheden invullen in deze formule

$$N \frac{m^2}{kg \times kg} = kg \times m \times s^{-2} \frac{m^2}{kg \times kg} = m^3 kg^{-1} s^{-2}$$

Hierin is gebruik gemaakt van het feit dat een newton ter herleiden is tot basiseenheden kg, m en s.

Hoofdstuk 32

1. Schattingen als deze zijn subjectief, maar dat neemt niet weg dat er wel wat over te zeggen is. De vraag is eigenlijk: wat is de “vroegste” tijd die je kunt aflezen en wat is de “laatste” tijd. De helft van het verschil kun je als de onzekerheid beschouwen. Anders gezegd: hoe goed kun je schatten tussen twee strepen op de klok. Zelf denk ik dat een kwart of een vijfde wel redelijk is. De afstand tussen twee strepen staat (wat de kleine wijzer betreft) voor één zestigste deel van twaalf uur en dus voor vijf minuten. Een onzekerheid van een minuut zou dan redelijk zijn. Ondertussen is het nog maar de vraag of die wijzer wel zo exact aangeeft hoe laat het is! Is de wijzerplaat exact rond? (Nee.) Zit het draaipunt van de wijzer exact in het midden? (Nee.) Zit er een beetje speling in het mechanisme dat de wijzer aandrijft? (Ja.)

De vraag was hoe groot de onzekerheid is *ten gevolge van het aflezen*. Die is ongeveer een minuut.
2. Dat is een uniforme verdeling die loopt tussen twee grenzen: de ondergrens is het aangegeven tijdstip in hele minuten, de bovengrens is één minuut later dan dat.
3. Tsja... Hoeveel is één gedeeld door vijf voor twee? Je zult de tijd moeten uitdrukken in minuten (of seconden) als eenheid. Seconden zijn “netter” omdat de seconde een SI-eenheid is. Minuten liggen meer voor de hand, gezien de aanduiding op de klok.

Langs de horizontale as staat de eenheid minuut.
Langs de verticale as staat de eenheid 1/minuut (of minuut^{-1}).

Hoofdstuk 34

1. Ook bij een analoge klok is het aantal cijfers of decimalen eindig. Het laatste cijfer kun je schatten, maar daarna komt er niet *nóg* een cijfer. De mogelijke waarden die je kunt aflezen zijn dus discreet.
2. Of je kunt schatten tussen de secondestreepjes hangt af van de vraag of de secondewijzer continu loopt of met stapjes verspringt.

Hoofdstuk 35

1. De onzekerheid bedraagt 0,75 V. De werkelijke waarde ligt namelijk tussen 482,5 V en 483,9999... V.
2. De spanning is $483,25 \pm 0,75$ V, wat in het juiste aantal significante cijfers gelijk is aan $483,3 \pm 0,8$ V.

Hoofdstuk 36

1. De relatieve onzekerheid van beide weerstanden samen (in serie) is ook 5%. De absolute onzekerheid is twee keer zo groot.
2. Reken alles om naar ml (cm^3) en absolute maten.
 $8\% \text{ van } 2400 \text{ ml} = 192 \text{ ml}$
 $2400 \text{ ml} + 600 \text{ ml} = 3000 \text{ ml}$
Als de relatieve onzekerheid in het totaal maximaal 10% mag zijn, mag de absolute onzekerheid in het totaal maximaal 300 ml zijn.
De absolute onzekerheden van beide ‘delen’ worden bij elkaar opgeteld. In het eerste deel zit al een onzekerheid van 192 ml. In het tweede deel mag dus nog een onzekerheid van 108 ml zitten.
Eigenlijk moet je dat schrijven in twee significante cijfers.
Dan wordt het 0,11 liter.
3. Er zijn vast een heleboel manieren, maar wat ik zelf zou doen, is: kopieer de tekst naar *Word* en voer de functie *Zoeken en vervangen* uit. Vervang dan de ‘m’ door een ander teken (of toch een ‘m’, dat mag ook) en *Word* vertelt je hoeveel keer hij die letter vervangen (en dus gevonden heeft.)
4. Een aantal aspecten speelt een rol.
 - a. Het *definitieprobleem*. Wat bedoel je met “hierboven”?
De alinea? De regel? De tekst van het hoofdstuk (of de pagina of het boek) tot aan dat punt?
Bedoel je alleen de kleine letter ‘m’ of ook de hoofdletter ‘M’?
 - b. *Met en Weten* heeft veel te maken met *nadenken* hoe je tot een uitkomst of waarde komt. Ook van zoiets doms als het tellen van de letter ‘m’.
Of het rare vraag is? Tsja...

Hoofdstuk 37

1. De absolute onzekerheid van een *verschil* tussen twee grootheden is gelijk aan *de som* van de absolute onzekerheden. Een relatieve onzekerheid van 5% op een

weerstand van $1,0 \text{ k}\Omega = 1000 \text{ }\Omega$ betekent een absolute onzekerheid van $50 \text{ }\Omega$. Twee keer die waarde is: $100 \text{ }\Omega$.

2. De vraag is in feite niet te beantwoorden. Het verschil tussen $1,0 \text{ k}\Omega$ en $1,0 \text{ k}\Omega$ is $0,0 \text{ k}\Omega = 0 \text{ }\Omega$. De absolute onzekerheid = $100 \text{ }\Omega$. Het verschil zou je dus moeten noteren als $0 \pm 100 \text{ }\Omega$.

Eigenlijk zeggen we dus: de weerstanden zijn aan elkaar gelijk, maar ze kunnen $100 \text{ }\Omega$ van elkaar verschillen.

De relatieve onzekerheid is rekenkundig gezien ongedefinieerd. Je zou ook kunnen zeggen: die is oneindig groot.

Hoofdstuk 38

1. Dit is een beetje een gekke vraag. Er staat namelijk niet in hoe groot de tafel is. Geen harde afmeting in meters, maar ook geen variabele.

De afstand tussen de poten is l , de hoogte van de kortste poot is h . Het hoogteverschil is d . We hoeven niet met hoeken (in graden of radialen) te rekenen. De verhouding overstaande gedeeld door aanliggende zijde is genoeg. Die is namelijk $3 \text{ mm} / 1 \text{ m} = 0,003$, gelijk aan de gegeven tolerantie van de waterpas. Dat is dezelfde verhouding als $d / L = 0,003$.gevraagd wordt naar d . Het antwoord is dus: $0,003 \cdot L$.

De tafel staat “waterpas” als het maximale hoogteverschil tussen de poten maximaal de tolerantie maal de lengte van het tafelblad is.

Let op: we doen nu net alsof een tafel een tweedimensionaal ding is. We veronderstellen dat hij vier poten heeft (wat niet eens gegeven was, drie of zes zou ook kunnen).

Hoofdstuk 39

Er is in dit sommetje sprake van alleen maar vermenigvuldigingen. Het volume van het stalen vat bedraagt $\text{lengte} \times \text{breedte} \times \text{hoogte}$ en het volume van het water erin vind je door het te vermenigvuldigen met een factor.

Eigenlijk klopt de vraag niet. Wat bedoel je met “de afmetingen” van het vat? De buitenkant of de binnenkant? Als het de buitenkant is, moeten de dikte van de wand(en) nog weten.

De relatieve fout wordt gevonden door de relatieve fout in die vier factoren op te tellen. Het zijn er vier omdat de onzekerheid in de hoogte is opgegeven als *fractie* van de hoogte. Was dat een *absolute* onzekerheid geweest, dan zouden er slechts *drie* factoren van belang zijn geweest.

De relatieve onzekerheden zijn $0,06 / 1 = 6\%$ en $0,06 / 4 = 1,5\%$ en $0,06 / 8 = 0,75\%$ en de gegeven 5% . De grootste bijdrage wordt dus geleverd door de 6% in de hoogte.

Als voor het hoogtepil van het water in plaats van de *fractie* van de hoogte van het vat een *absolute onzekerheid* van 5 cm was opgegeven, dan zou de hoogte niet meer ter zake hebben gedaan.

1. De relatieve onzekerheid van 6% in de hoogte van het vat levert de grootste bijdrage. Het is de grootste term van de vier bijdragen.
2. Dit is een gemene vraag. Want wat is de onzekerheid in het hoogtepeil van het water? Is die 5%? Nee. Althans: je zou de vraag kunnen lezen als: “De onzekerheid in de hoogte is 5% van exact 1 m.” Maar je kunt met evenveel recht beweren dat die onzekerheid gelijk is aan 5% van $(1,00 \pm 0,06)$ m. Dan is de onzekerheid dus een gevolg van die 5% én die 6%... Het gaat dan om het product van twee factoren met een onzekerheid. De totale relatieve onzekerheid is dan $5\% + 6\% = 11\%$.

Laten we de makkelijkste variant nemen. We bedoelen met 5% gewoon 5% van 1 m. Dat is (absoluut gezien) 5 cm.

Hoe zit dat nu met de vorige vraag? Hebben we die dan niet goed beantwoord? Dat ligt er ook weer aan hoe je het bekijkt. Die 5% als onzekerheid op “de helft” is niet de grootste bijdrage. Maar als je zegt dat de hoogte van het waterpeil gelijk is aan $0,50 \pm 0,06$ m, dan hebben we daarin de grootste bijdrage.

3. We beginnen altijd eerst met de onzekerheid. In de meest eenvoudige interpretatie van de vraag is de relatieve onzekerheid gelijk aan $6\% + 1,5\% + 0,75\% + 5\% = 13,25\%$.

Het volume is:

- lengte $\times$ breedte $\times$ hoogte $\times$ fractie =
 - $1 \text{ m} \times 4 \text{ m} \times 8 \text{ m} \times 0,5 =$
 - 16 m^3 .

Let (zoals altijd) op de eenheden. Lengte, breedte en hoogte zijn in meters. De fractie is dimensieloos.

De absolute onzekerheid is $13,25\%$ van $16 \text{ m}^3 = 2,12 \text{ m}^3$. In het juiste aantal significante cijfers is het volume van het water dus $16 \pm 2 \text{ m}^3$.

Het werd niet expliciet gevraagd, maar dat hoeft ook niet: de onzekerheid moet ALTIJD worden opgegeven als je die weet.

4. Het volume van een bol wordt gegeven door onderstaande formule.

$$V = \frac{4}{3}\pi r^3$$

De relatieve onzekerheid in het volume is drie keer de relatieve onzekerheid in de straal. Die mag dus 2% zijn: één derde van 6%.

De straal van de bol is 0,0500 m (de helft van de diameter van 1,00 dm). De absolute onzekerheid mag dus $6\% \times 0,0500 \text{ m} = 0,0030 \text{ m}$ zijn. In decimeters: 3,0 dm.

Hoofdstuk 40

```
som = 0
aantal = 0
for steen in range(1, 7):
    product = steen * steen
    som = som + product
    aantal = aantal + 1
gemiddelde = som / aantal
```

```

print(gemiddelde)

som = 0
aantal = 0
for steen1 in range(1, 7):
    for steen2 in range(1, 7):
        product = steen1 * steen2
        som = som + product
        aantal = aantal + 1
gemiddelde = som / aantal

print(gemiddelde)

```

Hoofdstuk 43

1. De formule voor optellen zegt dat je de absolute fouten bij elkaar optelt. En $\delta B + \delta B = 2\delta B$.

Hoofdstuk 44

1. Voor $n = 0$ heb je een grootte tot de macht 0. De waarde daarvan is 1. Omdat het exact 1 is, is de absolute onzekerheid gelijk aan 0 en de relatieve onzekerheid dus ook. Die is namelijk 0 gedeeld door 1.
2. Noteer eerst de formules voor de volumes van een bol en een kubus.

$$V_{bol} = \frac{4}{3}\pi x^3$$

$$V_{kubus} = x^3$$

Uit de derde macht in de formules blijkt dat de relatieve onzekerheid van beide volumes hetzelfde is, namelijk drie keer de relatieve onzekerheid van x .

$$\frac{\delta V_{bol}}{|V_{bol}|} = \frac{\delta V_{kubus}}{|V_{kubus}|} = 3 \frac{\delta x}{|x|}$$

De absolute onzekerheden zijn gelijk aan de relatieve onzekerheden maal de volumes zelf.

$$\delta V_{bol} = 3 \frac{\delta x}{|x|} \times V_{bol} = 3 \times \frac{4}{3}\pi x^3 \times \frac{\delta x}{|x|} = 4\pi x^2 \delta x$$

$$\delta V_{kubus} = 3 \frac{\delta x}{|x|} \times V_{kubus} = 3 \times x^3 \times \frac{\delta x}{|x|} = 3x^2 \delta x$$

Om die twee uitkomsten met elkaar te vergelijken, is het handig om ze op elkaar te delen.

$$\frac{\delta V_{bol}}{\delta V_{kubus}} = \frac{4\pi x^2 \delta x}{3x^2 \delta x} = \frac{4}{3}\pi \approx 4,19$$

De absolute onzekerheid in het volume de bol is dus ruim vier keer zo groot als de absolute onzekerheid in het volume de kubus.

3. Het maakt niet uit of je naar de straal of de diameter kijkt, omdat in beide formules een kwadraat staat. Beide onzekerheden worden een factor vier groter.

Hoofdstuk 48

Als je een reeks gehele getallen met elkaar vermenigvuldigt en er zit een (veelvoud van) tien bij, dan eindigt het product altijd op een 0. Waar je verder ook mee vermenigvuldigt: dat laatste cijfer zal altijd zo blijven.

Als je een reeks gehele getallen met elkaar vermenigvuldigt en er zit minstens één even getal tussen, dan is de uitkomst altijd even. Waar je verder ook mee vermenigvuldigt: het getal zal altijd even blijven.

In de reeks 1 t/m 9 zit geen 10, maar wel een 5. En als je 5 met een even getal vermenigvuldigt, komt er een veelvoud van 10 uit. Met een 0 op het eind dus.

Bovenstaande redeneringen maken het aannemelijk dat de 0 op het einde vaak zal voorkomen. De kans dat er geen enkel even getal in de reeks zit, is nogal klein. En de kans dat er minstens één 5 voorkomt, is juist niet klein. Die is namelijk $1 - (8/9)^9 \approx 65,4\%$. De kans alle negen cijfers oneven zijn, is $(5/9)^9 \approx 0,50\%$. Als er één 5 bij zat, is die kans al groter. Dan is $(5/9)^8 \approx 0,90\%$.

De kans op een oneven uitkomst is wél klein. Die heb je alleen maar als alle negen cijfers oneven waren. (Als alleen 1, 3, 5, 7 en 9 voorkomen dus.)

De kans op een 0 aan het einde is groter dan alle andere kansen samen. De kans op een oneven getal is erg klein, dus de 1, 3, 5, 7 en 9 zullen niet hoog scoren. Als we ervan uitgaan dat de even getallen even waarschijnlijk zijn, is de kans op 2, 4, 6 en 8 dus ongeveer $34,6\% / 4 = 8,65\%$.

Mijn prognose van het histogram is (niet-grafisch) weergegeven in de volgende reeks.

[651, 1, 86, 1, 86, 1, 86, 1, 86, 1]

(Het totaal is 1000 en ik heb gewerkt met gehele getallen.)

Hoofdstuk 58

1. Je kunt het zeggen, maar er is nogal wat op aan te merken.
 - a. Om met de belangrijkste te beginnen: een Celsiusschaal is geen ratioschaal. We vergelijken de waarden 293,15 K en 293,35 K met elkaar. Dat verschilt niet 1%, maar slechts 0,07%.
 - b. De gegevens zijn onzorgvuldig en onvolledig geformuleerd. Er wordt niet eens vermeld of de meting op hetzelfde tijdstip is uitgevoerd. Het kan bovendien zo zijn dat de minimumtemperatuur *lager* is en de maximumtemperatuur *hoger*. Bij helder weer is dat mogelijk. Vaak kijken we naar het maximum, maar het gemiddelde ligt meer voor de hand. Dat het 'vandaag' warmer is, kun je ook zeggen na een heel koude nacht.
 - c. Of het *warmer* is, is niet alleen een kwestie van temperatuur, maar ook van hoeveel zonneschijn er is, de windkracht, luchtdruk en luchtvochtigheid.
 - d. Er is geen uitspraak gedaan over de onzekerheid in de genoemde temperaturen. Op een huis-, tuin- en keukenthermometer zal een verschil van 0,2 °C niet significant zijn.
2. In symbolen kun je zeggen: $A > B$, $B > C$, $C > A$.
Wiskundig gezien kan dat niet. Als A groter is dan B en B is groter dan C, dan moet A groter zijn dan C.

Er is bij een schaakwedstrijd sprake van een *ordinale* schaal. Je kunt winnen of verliezen, de wedstrijd kan ook in remise eindigen: gelijk. Het probleem met deze ordinale schaal is dat die alleen maar geldt *voor die ene wedstrijd*. Je kunt niet zomaar zeggen dat de ene schaker ‘beter’ is dan de andere. Tenzij die ene altijd wint.

3. Een goede analyse begint met de harde cijfers.

Kees 37%, Piet 35%, Jan 27%.

In totaal is dat 99% en geen 100%.

Er zijn twee mogelijke oorzaken:

- a. De percentages zijn afrondingen van (bijvoorbeeld) 37,2%, 35,5% en 27,3%. De som daarvan is wel 100%.
- b. Er kunnen blanco stemmen zijn. Die tellen wel mee in het percentage (in tegenstelling tot ongeldige stemmen), maar worden niet uitgebracht op de drie kandidaten.

Tweede ronde:

Kees 41% (4% meer) en Jan 59% (32% meer)

Hoe kan het dat Piet in de tweede ronde veel meer stemmen kreeg dan in de eerste, terwijl dat verschil bij Kees zo klein is? Een mogelijke verklaring is dat de mensen die hun stem uitbrachten het liefst stemden op iemand uit een bepaalde plaats. Wie een voorkeur had voor een inwoner van Tilburg moest kiezen tussen twee kandidaten. Er waren mensen die op Piet stemden terwijl ze liever Jan hadden dan Kees. In de eerste ronde was dat niet zichtbaar, in de tweede wel.

De stemming was een verkapt verkiezing tussen Eindhoven en Tilburg. Een binaire en ordinale schaal liepen door elkaar heen.

Theoretisch is het mogelijk dat Jan in een tweede ronde óók gewonnen zou hebben als hij de enige kandidaat tegenover Kees was geweest. Sterker nog: misschien zou Jan wel gewonnen hebben als hij de enige kandidaat tegenover Piet was geweest!

Als alle Eindhovenaren gedacht zouden hebben: ‘Dan altijd nog liever Jan...’

Hoofdstuk 60

1. Het verschil is waarschijnlijk niet significant, omdat een tentamen een meting is die niet op één tiende punt nauwkeurig kan vaststellen hoe goed een student de stof beheerst. Wie een 5,4 haalde, had met een net iets andere vraag in de toets misschien wel een 5,7 gehaald. Of een 4,9.

Je kunt een tentamencijfer beschouwen als een experiment met een normale verdeling. Het gemiddelde van die verdeling is hoger bij iemand die de stof beter beheerst. De standaardafwijking is aanzienlijk groter dan 0,1 punt. Hoe groot die is, hangt af van het vak en de docent en de manier van toetsen.

2. Het verschil is waarschijnlijk wel relevant. Het is namelijk het verschil tussen sla-gen en zakken.

Eigenlijk kun je zonder de context de vraag niet beantwoorden. Voor iemand die besluit om te stoppen met de studie of gewoon nooit af te studeren geeft het niet dat het cijfer onvoldoende is.

Hoofdstuk 62

1. Aanpak A, want die kost minder moeite.

2. Aanpak B, want die verlaagt het uitvalspercentage (iets) verder.
Je kunt ook antwoorden: dat weet je niet, want een verschil van 1% zal niet significant zijn. En is ook nauwelijks relevant.
3. Er wordt alleen naar percentages gekeken en niet naar absolute aantallen. Als de instroom als geheel sterk terugloopt door die toelatingseis, betaal je wel een hoge prijs voor het verlagen van de uitval. Als er van 200 studenten 90 overblijven, is dat wellicht wenselijker dan dat er van 150 studenten 63 overblijven.
4. Ja, een methode kan zowel het doel beter bereiken als minder moeite kosten. (De NBSA-norm naar 10 studiepunten verlagen bijvoorbeeld.)

Hoofdstuk 63

1. Zie [ISO 25010 - Wikipedia](#)
2. 31
3. 25010

Hoofdstuk 64

1. $1,05 \times 0,5\% = 0,525\%$.
Uitgedrukt in één significant cijfer bevat het nog steeds (maximaal) 0,5% alcohol.
2. Maximaal 0,5% betekent minimaal 0,0% en maximaal 0,5%.
Bij beide grenzen moet je 5 pp optellen.
Het bevat dus minimaal 5,0% en maximaal 5,5% alcohol.
3. $174 / 255 = 0,682352941176471$ volgens *Excel*.
Maar ja, is dat het zo groot mogelijke aantal decimalen? Voor *Excel* wel. De [Full Precision Calculator \(mathsisfun.com\)](#) zou het liefst met volledige precisie werken. *But some decimals go on forever, so we stop after 200 decimals*. Dit is het getal met tweehonderd cijfers achter de komma:
0,6823529411764705882352941176470588235294117647058823

5294117647058823529411764705882352941176470588235294117647058823529411764
7058823529411764705882352941176470588235294117647058823529411764705882352
94.

Het (b)lijkt inderdaad een repeterende breuk te zijn. Als je spaties zet tussen de repeterende “blokken”, dan staat er namelijk:
0,6 8235294117647058 8235294117647058 8235294117647058 ...

Antwoord:

RGB(0%, 0%, 0,6823529411764705882352941176470588235294117647%)

Het zo groot mogelijk aantal dat nog op één regel past.
4. Volgens *Excel*:
=NORM.VERD.N(1;0;1;WAAR) - NORM.VERD.N(-1;0;1;WAAR)
0,682689492137086
5. Als je beide waarden in *Excel* onder elkaar zet, kun je ze “handmatig” van elkaar aftrekken. We zetten het hoogste getal boven.

$$\begin{array}{r} 0,682689492137086 \\ - 0,682352941176471 \\ \hline 0,000336550960615 \end{array}$$

Als je *Excel* de waarden van elkaar laat aftrekken, vind je iets anders. In vijftien decimalen is het 0,000336550960616. Als je de berekening in twintig decimalen probeert uit te voeren, kom je op het volgende.

$$\begin{array}{r} 0,68268949213708600000 \\ - 0,68235294117647000000 \\ \hline 0,00033655096061591100 \end{array}$$

Als je twee getallen met vijftien decimalen van elkaar wilt aftrekken en je wilt het exacte antwoord weten, dan kun je het beter met pen en papier uitrekenen dan met *Excel*. Tenzij je het wat sneller wilt doen en een verschil van 0,000000000000001 (10⁻¹⁵) acceptabel vindt.

Nu nog het antwoord in procentpunt: 0,000336550960615 pp.

6. 0,034%

Hoofdstuk 65

1. De getallen (zonder de eenheden) worden beide door 2,54 gedeeld, dus de relatieve fout verandert niet.
2. Dit is een beetje een strikvraag, maar hij klopt wel.

Ook de absolute fout verandert niet! De *getallen* veranderen wel (en de absolute fout dus ook), maar de *grootheid zelf* verandert niet.

Voorbeeld: de tafel is 2,5400 m lang en de onzekerheid is 1,27 cm. (Het aantal significante cijfers in de onzekerheid is in dit voorbeeld voor het gemak even te groot gekozen.) Als je dat omrekent naar inches, kom je op een lengte van 100,0 inch en een onzekerheid van 0,5 inch. Het getal 0,5 is natuurlijk wel kleiner dan het getal 1,27, maar 0,5 inch is niet kleiner dan 1,27 cm. Het is precies hetzelfde.

3. Zowel voor de flitspalen als voor de trajectcontrole geldt dat het maar een zeer beperkte meting is. De flitspaal meet maar op één moment, dus de auto kan *daarvoor en daarna* te hard rijden.

Voor de trajectcontrole geldt dat de onzekerheid veel kleiner is, maar ook daar weet je niets over ervoor en erna. Ook op het de traject waar gemeten is, kan de snelheid op een bepaald moment te hoog geweest zijn. Je weet alleen het *gemiddelde* op dat traject.

Hoofdstuk 70

1. Geen uitwerking: controleer of je spreadsheet klopt.
2. =GEMIDDELDE(A:A)
Deze formule bepaalt het gemiddelde van een hele kolom. Daardoor is voldaan aan de eis dat het aantal elementen in kolom A veranderd kan worden.
3. De *Excel*-formules en waarden staan in onderstaande tabel.

=GEMIDDELDE(A:A)	23
=SOM(A:A)	322
=AANTALARG(A:A)	14
=MEDIAAN(A:A)	23
=MODUS.ENKELV(A:A)	23
=MAX(A:A)	24
=MIN(A:A)	21
=ALS(B6-B1>B1-B7;B6-B1;B1-B7)	2
=GEM.DEVIATIE(A:A)	0,714285714
=AANTAL.ALS(A:A;">"&B1)	5

4. De veranderingen staan in onderstaande tabel.

gemiddelde	blijft 23
som	wordt 345 (23 meer)
aantal waarden	wordt 15 (1 meer)
mediaan	blijft 23
modus	blijft 23
grootste waarde	blijft 24
kleinste waarde	blijft 21
grootste afwijking van het gemiddelde	blijft 2
gemiddelde absolute afwijking	wordt 0,666667
aantal waarden groter dan gemiddelde	blijft 5

5. Het gemiddelde, de som, de kleinste waarde en de gemiddelde absolute afwijking veranderen.
 Eén waarde verandert spectaculair: het aantal waarden groter dan het gemiddelde wordt “opeens” gelijk aan 11, terwijl het 5 was. Dat lijkt vreemd, maar het gemiddelde wordt een klein beetje kleiner. Daardoor zijn de zes waarden van 23 “opeens” groter dan het gemiddelde van 22,8.
6. Voer zelf de test uit.
7. Als de kans dat iets gebeurt P is, is de kans dat het *niet* gebeurt $1 - P$. De kans dat *iemand* iets heeft of doet, bereken je door eerst te bepalen hoe groot dat kans is dat *niemand* iets heeft of doet. Die kans is $(1 - P)^N$ als N het aantal personen is.
 Als we rekenen met $P = 10^{-9}$ en $N = 8 \cdot 10^9$, vinden we dat de kans dat niemand iets heeft of doet maar 0,034% is. (Als we hadden gerekend met één miljard mensen, zouden we op ongeveer 37% zijn uitgekomen.)
 De kans dat iemand het heeft of doet is dus ruim 99,9%.

Hoofdstuk 71

Je zou een byte kunnen splitsen en de eerste vier bits gebruiken voor het cijfer voor de komma en het laatste deel van vier bits voor het cijfer na de komma. Met vier bits kun je zestien cijfers aan. Je hebt er tien nodig voor de komma (1 t/m 10) en tien na de komma (0 t/m 9).

Hoofdstuk 73

1. Er zijn 36 mogelijkheden met een gelijke kans. Bereken het kwadraat van elk van die mogelijkheden en bepaal het gemiddelde. Dat geeft (in twee decimalen) 54,33.

De verwachtingswaarde van het totaal per worp is overigens 7.

Het kwadraat van de verwachtingswaarde is dus niet gelijk aan de verwachtingswaarde van het kwadraat.

Om over na te denken: de verwachtingswaarde bij één dobbelsteen is $3\frac{1}{2}$, maar je kunt nooit $3\frac{1}{2}$ gooien.

2. Het klopt dat in de ene envelop het dubbele zit van het bedrag van de andere. Maar dat is dan wel het dubbele *van een ander bedrag*. In de ene envelop zit X euro en in de andere 2X euro. Je weet niet welke van de twee je hebt. Heb je die met 2X, dan verlies je X door te ruilen. Heb je die met X, dan win je X door te ruilen. In de ‘hoogste’ envelop zit het dubbele van X. In de ‘laagste’ zit de helft van 2X. Je wint dus X óf de helft van 2X.
3. De verwachtingswaarde is gelijk.
In beide gevallen het dubbele van één keer meedoen met één loterij.
4. De vraag is niet te beantwoorden, omdat hij niet goed gedefinieerd is. ‘Bij vijftig loten’ kan op twee manieren worden geïnterpreteerd:
 - a. als je vijftig loten *koopt* (en niet één of twee)
 - b. als er vijftig loten *te koop zijn* (en niet honderd)
5. Bij 0 en 100: geen enkel lot, of alle loten. In beide gevallen ligt de uitkomst vast en is de winst geen kansvariabele meer.
6. De verwachtingswaarde van de winst per lot is vast. *De verwachtingswaarde van de totale winst* hangt dus af van het aantal lootjes dat je koopt. Twee keer tachtig is meer dan 92.

De kans dat je een prijs wint, is een ander verhaal. Als je 92 lootjes koopt, ben je zeker van een prijs. Er zijn namelijk $1 + 3 + 5 = 9$ lootjes waarop een prijs valt. Het kan niet zo zijn dat je die alle negen niet hebt, want je hebt er 92 van de 100. Als je extreem veel pech hebt, geef je € 92 uit aan lootjes en heb je één prijs van € 5.

Bij twee keer tachtig lootjes is het mogelijk dat je geen prijs hebt.

7. Argumenten om wel of niet mee te doen zijn heel subjectief. Wie vindt dat de naam *Nationale Postcode Loterij* fout gespeld is (postcode-loterij is één woord), heeft gelijk. Is dat een argument om niet mee te doen? Voor sommige mensen wel. Wie principieel tegen gokken is, doet sowieso niet mee. Wie zich geen raad zou weten met een grote prijs, ook niet. Wie het ‘gewoon leuk’ vindt, heeft een heel eenvoudige reden. Het steunen van een goed doel kan ook een reden zijn.

Rationeel gezien kan een argument zijn: de kosten van een lot zijn niet heel hoog. Dat geld kun je wel missen. Dat je een heel groot geldbedrag kunt winnen, is een aantrekkelijke gedachte. Dat de winstverwachting negatief is, wil op zichzelf niet zeggen dat je het niet moet doen.

Hoofdstuk 79

1. Bereken wanneer je die waarde exact bereikt.

Het juiste antwoord moet nog wel het woord *minimaal* bevatten. Je moet namelijk *minimaal* 400 metingen doen. Meer is ook goed.

- $$\frac{1}{\sqrt{2N-2}} = 0,1 \Rightarrow \sqrt{2N-2} = 10 \Rightarrow 2N-2 = 100 \Rightarrow N = 51$$

Het juiste antwoord moet nog wel het woord *minimaal* bevatten. Je moet namelijk *minimaal* 51 metingen doen. Meer is ook goed.

- “In het algemeen” ligt $\bar{x}$ dicht bij μ , in die zin dat de spreiding van een gemiddelde van tien kleiner is dan de spreiding van één waarde. De standaardafwijking is ruim drie keer zo klein. Maar het is niet *algemeen* waar dat één meting verder van het gemiddelde ligt dan het gemiddelde van een steekproef.

Het gemiddelde van deze tien metingen is 8,03. (De steekproef-standaardafwijking is 0,894, maar doet niet ter zake.) Het gemiddelde van de metingen wijkt (in dit specifieke geval!) drie keer zoveel af van het gemiddelde van de verdeling als de eerste meting. Ja, dat is toeval. En meestal is het andersom. Je kunt dus niet zeggen dat $\bar{x}$ “in het algemeen” dichter bij μ ligt dan die ene meting.

[illegible]

$$\frac{\delta l}{|l|} = 2 \frac{\delta T_0}{|T_0|}$$

De slinger is 3,4018 m met een onzekerheid van 0,3676 m.

In het juiste aantal cijfers: $3,4 \pm 0,4$ m.

In plaats van de onzekerheidsformule, kun je ook de hoogste en laagste waarde invullen die horen bij $3,7 \pm 0,2$ s, namelijk 3,5 s en 3,9 s. De lengtes die je dan vindt zijn 3,044 m en 3,780 m. Je komt met die aanpak op 3,4118 m met een onzekerheid van 0,3678 m. In het juiste aantal cijfers: $3,4 \pm 0,4$ m.

Hoofdstuk 86

Het handigste is om met een standaardnormale verdeling te rekenen.

```
import numpy as np
import statistics
from scipy.stats import norm
import time

aantal_metingen = 10
aantal_herhalingen = 10000
totaal = 0
aantal_verworpen = 0

begintijd = time.time()
for _ in range(aantal_herhalingen):
    # genereer een aantal random getallen
    metingen = np.random.normal(size=aantal_metingen)

    # bepaal de steekproef-standaardafwijking
    std_afw = statistics.stdev(metingen)
    gemiddelde = statistics.mean(metingen)

    # bepaal de grootste afwijking
    grootste = max(abs(metingen - gemiddelde))

    # bepaal de overschrijdingskans en de verwachtingswaarde
    overschrijdingskans = (1 - norm.cdf(grootste, gemiddelde, std_afw)) * 2
    verwachtingswaarde = overschrijdingskans * aantal_metingen

    totaal = totaal + verwachtingswaarde

    if verwachtingswaarde < 0.5:
        aantal_verworpen = aantal_verworpen + 1

# druk het percentage af waarbij de grootste waarde verworpen is
print("percentage verworpen: ", round(aantal_verworpen / aantal_herhalingen * 100, 1))
print("rekentijd (seconden): ", round(time.time() - begintijd, 1))
```

Hoofdstuk 88

Er zijn er genoeg te vinden.

Hoofdstuk 89

	A	B	C	D	
1	3,8	gemiddelde	3,38	3,383333333	=GEMIDDELDE(A:A)
2	3,5	hoogste	3,9	3,9	=MAX(A:A)
3	3,9	laagste	1,8	3,9	=MIN(A:A)
4	3,9	afwijking 1	0,52	1,8	=C2-C1
5	3,4	afwijking 2	1,58	0,516666667	=C1-C3
6	1,8	grootste afwijking	1,58	1,583333333	=MAX(C4:C5)
7		"rough and ready"	1,06	1,055555556	=2*C6/3
8					
9		mediaan	3,65	3,65	=MEDIAAN(A:A)
10		afwijking 1	0,25	0,25	=C2-C9
11		afwijking 2	1,85	1,85	=C9-C3
12		grootste afwijking	1,85	1,85	=MAX(C10:C11)

Hoofdstuk 94

1. Niet correct, maar leuk geprobeerd, is onderstaande oplossing.

```
a = 1
b = 2
c = (a + b)/10
print(c)
```

Het lijkt te werken, maar dat is meer geluk dan wijsheid. Deze c kan namelijk niet exact aan 0,3 gelijk zijn, aangezien dat niet te representeren is met een float. Dit is hoe het wel moet.

```
import decimal
from decimal import Decimal

a = Decimal('0.1')
b = Decimal('0.2')
c = a + b
print(c)
```

2. Rekenen met breuken: maak ze gelijknamig, in dit geval met 21 in de noemer.

$$\frac{2}{3} + \frac{3}{7} = \frac{2 \times 7}{3 \times 7} + \frac{3 \times 3}{7 \times 3} = \frac{14}{21} + \frac{9}{21} = \frac{14 + 9}{21} = \frac{23}{21} = 2 \frac{2}{21}$$

3. Gebruik de module fractions

```
from fractions import Fraction

a = Fraction(2,3)
b = Fraction(3,7)
c = (a + b)
print(c)

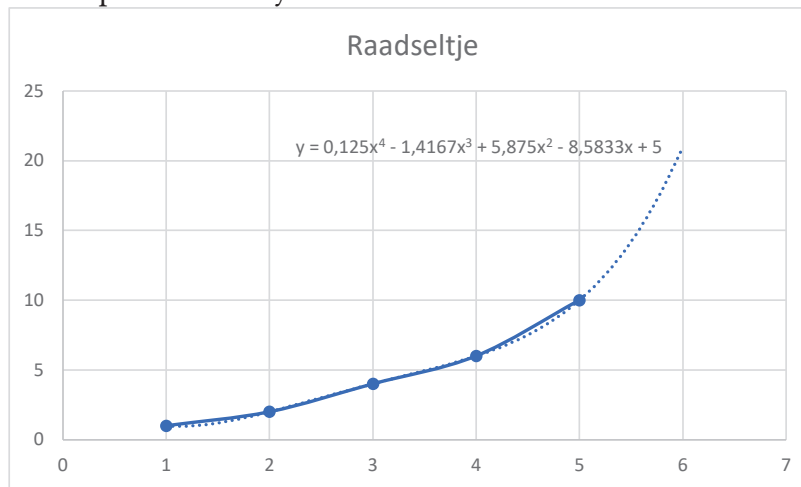
# Het kan ook in één regel
print (Fraction(2,3) + Fraction(3,7))
```

Hoofdstuk 98

- a) Te vinden met Google en de juiste zoektermen.

```
import numpy as np
x = np.array([1, 2, 3, 4, 5])
y = np.array([1, 2, 4, 6, 10])
z = np.polyfit(x, y, 4)
print(z)
```

- b) Zet de data onder elkaar in een kolom in *Excel*. Selecteer deze data en kies op de tab Invoegen in de groep Grafieken voor een spreidingsdiagram. Selecteer de grafiek en klik op Grafiekelementen (het kruis aan de rechterkant). Kies dan voor Trendlijn | Meer opties... en Polynoom.



- c) Het narekenen kan via een tabel in *Excel*.
In de linker kolom staan de getallen 1 t/m 6. In de bovenste rij staan de termen uit de formule. In de rechter kolom staat de som van de kolommen 2 t/m 6. In de kolommen 2 t/m 6 staan de termen die bijdragen aan de som, via bijvoorbeeld de formule $=D3*D3*D3*E\$1$.

	5	-8,5833	5,875	-1,4167	0,125	
1	5	-8,5833	5,875	-1,4167	0,125	1
2	5	-17,1666	23,5	-11,3336	2	2
3	5	-25,7499	52,875	-38,2509	10,125	4
4	5	-34,3332	94	-90,6688	32	6
5	5	-42,9165	146,875	-177,0875	78,125	10
6	5	-51,4998	211,5	-306,0072	162	21

Uitleg: D3 wordt met zichzelf vermenigvuldigd om de juiste macht te krijgen (in dit geval 3). E duidt op de vijfde kolom (waar de term voor derde macht staat). Het dollartekentje is er om ervoor te zorgen dat de formule gekopieerd kan worden naar andere velden.

Bij het uitrekenen vind je niet de waarden 1, 2, 4, 6 en 10, maar 1, 1,9998, 3,9992, 5,998 en 9,996. De oplossing is ook niet exact. Voeg je in de termen $-8,5833$ zo veel

mogelijk drieën toe en maak je van $-1,4167$ het getal $-1,41666666666667$ (zo veel mogelijk zessen en de laatste een 7) dan kom je in vijftien decimalen op deze uitkomsten.

1,0000000000000000
1,999999999999980
3,999999999999910
5,999999999999800
9,999999999999570

Niet exact dus, maar dat lukt nu eenmaal niet met een eindig aantal decimalen.

2. De reeks 1, 2, 4, 6, 10 wordt gevormd door het getal 1 gevolgd door $1+1$. Het tweede getal is dus eentje meer dan het eerste getal. Daarna komen er twee getallen die 2 groter zijn dan het vorige en daarna eentje die 4 groter is dan de vorige.

De getallen worden dus verhoogd met 2^0 , 2^1 , 2^1 en 2^2 . Eerst één keer de macht 0, dan twee keer de macht 1... Daarna? Drie keer de macht 2? Of vier keer? Als je niks beters kunt bedenken zou je 14 kunnen zeggen.

Andere optie: kijk eens naar onderstaande zin.

O, en hier passen expliciete aandoenlijke aandachttrekkers: tijdsverspillingen, cryptogrammenmakertjes!

Het aantal letters van de woorden: 1, 2, 4, 6, 12, 16 en 22.

Wie lijdt aan hippopotomonstrosesquippedaliofobie, moet niet meer verder gaan.

Hoofdstuk 99

1. De coëfficiënten zijn nu genoteerd als breuken. Als je ze uitrekent als decimale getallen met een komma, zie je dat de eerste veel groter is dan de volgende. In de linker kolom staan de uitgerekende waarden, rechts de verhouding tot het eerste getal.

0,0625	1
0,00358	0,0573
0,000235	0,00375
0,000174	0,00278

Als je ook nog rekening houdt met het feit dat alle hoeken kleiner zijn dan 0,7 radianen, worden de verschillen een factor 2, 4 en 8 groter. Je verwaarloost dus maar een procent of drie ten opzichte van die eerste term.

Hoofdstuk 107

1. Hier is misschien wel sprake van een definitieprobleem. Tellen alle pagina's mee, of alleen de bedrukte? Waar begin je met nummeren en wat doe je met de omslag van het boek?

Dat aantal van 366 is gebaseerd op het aantal dat *Microsoft Word* zelf aangeeft.

- a. Met bovenstaande definitie is de stelling juist. (Tenzij er iets misgegaan is bij het aanmaken van de PDF.)

- b. De auteur weet natuurlijk niet wat de lezer denkt, maar wat hij gedaan heeft, is het volgende. In *Word* zit de mogelijkheid om een *veld* op te nemen dat het aantal pagina's toont. Als je daarvan gebruik maakt, hoef je zelf niet te tellen. Belangrijker dan dat: als het aantal pagina's in een volgende uitgave verandert, blijft het gewoon kloppen. In het menu *Invoegen* zit een knopje *Snelonderdelen* en daaronder *Veld*. In het lijstje waaruit je kunt kiezen, zit een veldnaam *NumPages*.
- c. Het aantal even pagina's is het aantal pagina's (366) gedeeld door 2, afgerond naar beneden. Het aantal woorden is (volgens het veld *NumWords*) gelijk aan 90206. (Hoe weten we of dat getal klopt?)

Aantal pagina's	366	<i>NumPages</i>
Aantal fouten	122	=ROUND(B1/3;0)
Aantal woorden	90206	<i>NumWords</i>
Percentage	0,13525%	=B2/B3*100

De vraag is hoe je met dat aantal fouten moet omgaan. Delen door 3 inderdaad, maar moet het een geheel getal zijn voordat je verder rekt? En rond je dat getal dan "gewoon" af of naar beneden? Het maakt voor het eindantwoord niet veel uit, maar het is wel het verschil tussen exact en (net) niet exact. Het verschil is minder dan één procent van nog geen twee promille.

Waarom heb je deze oefening eigenlijk gedaan? Wat was het nut, het *leerdoel* zoals docenten zeggen?

- Je hebt als het goed is geleerd om met velden te werken in plaats van met harde getallen.
- Je hebt gezien dat je in *Word* ook kunt rekenen met velden.
- Je bent je gaan realiseren dat een boek met weinig fouten (eentje per drie pagina's is toch best netjes) er toch al gauw meer dan honderd bevat.
- Je ziet in dat we met zo'n groot aantal nog steeds over promillen spreken.

2. Het gaat wat te ver om deze berekening uit te voeren met velden en die kostprijs zal jaarlijks veranderen. Laten we voor het gemak uitgaan van 320 pagina's, 25 euro en 100 grams papier. Het boek is grofweg A5-formaat. A5 is de helft van A4, een kwart van A3, een achtste van A2, een zestiende van A1, een tweeëndertigste van A0. Een A0 is één vierkante meter. Het boek is dus 10 m² papier van 100 g/m². Dan zou het binnenwerk dus 1 kg wegen. Als papier € 50 per ton kost, dan zit er dus voor 1/1000 van € 50 aan kosten in het papier. Dat is vijf eurocent...

Je kunt ook een ander sommetje maken. Een pak papier van 500 vel kost in de winkel een eurootje of zes. Je kunt een vel papier aan twee kanten bedrukken en A4'tjes doormidden snijden om er A5 van te maken. Met één pak papier druk je dus vier boeken van 500 pagina's. Dan zijn de kosten niet tien eurocent, maar € 1,50. Een kosten van het papier in een pak papier zijn niet zo heel hoog. Wat zit er in een pak papier nog meer dan alleen papier?

Hoofdstuk 100

1. $\sin(0,01) = 0,009999833$
Het verschil is 0,017‰.
2. $\sin(x) = x - \frac{1}{3!}x^3 + \frac{1}{5!}x^5 - \frac{1}{7!}x^7 + \dots$

$$\sin(0,001) = 0,001 - \frac{1}{6}10^{-9} + \frac{1}{120}10^{-15} - \frac{10^{-21}}{5040} + \dots$$

$$= 0,0010000001667$$

De tweede term is 1/6 van een miljoenste van de eerste. Daar kun je al nauwelijks meer rekening mee houden. Je hebt acht significante cijfers nodig om het verschil te noteren.

Hoofdstuk 101

1. Je kunt alleen maar vermenigvuldigen en delen als je met een absolute temperatuurschaal werkt. De temperaturen dienen daarom te worden omgerekend van graden Celsius naar kelvin. Daarna kun je ze op elkaar delen.

$$\frac{T_2}{T_1} = \frac{273,15 + 20}{273,15 + 10} = 1,03531697$$

Omdat beide temperaturen zonder onzekerheid zijn opgeven in twee significante cijfers, kiezen we in het antwoord ook voor twee significante cijfers.

Antwoord: 3,5%

2. Eigenlijk is de opgave niet compleet. Er zou nog bij moeten staan dat Jeroen wil werken met een lineaire schaal. Er zijn namelijk oneindig veel functies die door de twee gegeven punten lopen. Een lineair verband ligt wel het meest voor de hand.

$$T_J = aY_C + b$$

$$a = \frac{T_{J2} - T_{J1}}{T_{C2} - T_{C1}} = \frac{100 - 0}{0,01 - (-273,15)} = \frac{10000}{27316}$$

$$b = T_{J1} - aT_{C1} = 0 - \frac{100}{273,16} \cdot (-273,15) = \frac{2731500}{27316}$$

In de berekening van a en b dienen alle onafgeronde getallen te blijven staan. Afronden op (bijvoorbeeld) drie significante cijfers zou maken dat er geen exact antwoord meer gegeven wordt. Afronden doe je pas als je getallen invult met een eindige nauwkeurigheid. Bij het programmeren van de oplossing in *Python* of het invoeren van een formule in *Excel* laten we bij voorkeur de breuken staan, tenzij er redenen zijn om dat niet te doen. In dat geval luidt de formule als volgt.

$$T_J = 0,366 \cdot Y_C + 100$$

Het getal 100 is een afronding van 99,996. Omdat we met twee significante cijfers rekenen, zou deze formule bruikbaar zijn. Invullen van 20 °C resulteert in 107,32 °J. Als je strikt bent in het werken met twee significante cijfers, zou je dat moeten noteren als $1,1 \cdot 10^2$ °J. Verdedigbaar is om niet naar het aantal significante cijfers te kijken, maar in hele graden te rekenen.

Antwoord: 20 °C komt overeen met 107 °J.

Bijlagen

A. Alt-codes en Unicodes

Tekens die niet op het toetsenbord zitten, kun je in *Word* invoegen door de Alt-toets in te drukken en tegelijkertijd de ASCII-code van dat teken in te tikken op het numerieke toetsenbord. Als het teken niet in de ASCII-tabel zit, bestaat er vaak een code voor die met een 0 begint. Een alternatief is de Unicode intikken en daarna op Alt+X drukken. Als het teken niet los staat van een ander teken, moet je de code (het getal) eerst apart selecteren en dan op Alt+X drukken. Merk op dat de unicodes hexadecimale getallen zijn.

Het vaste streepje en de vaste spatie (streepjes en spaties die ervoor zorgen dat de delen die er links en rechts van staan op verschillende regels komen) kun je ook maken door Ctrl+Shift in te drukken en de spatie of het streepje in te tikken.

Een *vaste spatie* wordt ook wel een *harde spatie* genoemd.

De ‘belangrijkste’ twee uit deze tabel zijn het vermenigvuldigingsteken en het minteken. Bij dat minteken helpt *Word* soms een handje. Die maakt (afhankelijk van instellingen) in de som $9 - 3 = 4 \times 1\frac{1}{2}$ van de - een – en van $1/2$ maakt hij $\frac{1}{2}$. Maar van een * (sterretje) of een x (kleine letter X) maakt hij niet een ‘net’ vermenigvuldigingsteken.

Maar... van de – maakt hij geen ‘net’ minteken. *Word* neemt een *half kastlijntje* (Unicode 2213). Je ziet het verschil bijna niet. In het lettertype van *Meten & Weten* is het verschil 3% en de breedte en 1% in de hoogte.

In de tabel hieronder staan de breedte en hoogte volgens *Inkscape* bij *Linux Libertine 12*.

		Unicode	breedte	hoogte	
$9 - 3 = 6$	echt minteken	2212	1,821	0,211	
$9 - 3 = 6$	half kastlijntje	2013	1,881	0,209	Ctrl+‘-’
$9 - 3 = 6$	heel kastlijntje	2014	2,811	0,209	Ctrl+Alt+‘-’
$9 - 3 = 6$	koppelteken	002D	1,091	0,254	‘-’

	Alt	Unicode	omschrijving
	8201	2009	smalle spatie (<i>thin space</i>)
	0160	00A0	vaste spatie (<i>non-breaking space</i>)
-	8209	2011	vast streepje (<i>non-breaking hyphen</i>)
–	0150	2013	half kastlijntje (<i>en dash</i>)
—	0151	2014	heel kastlijntje (<i>em dash</i>)
–	8722	2212	min (<i>minus</i>)
×	0215	00D7	vermenigvuldigen
÷	0247	00F7	delen
±	241	00B1	plusminus
‰	0137	2030	promille
$\frac{1}{4}$	0188	00BC	kwart
$\frac{1}{2}$	0189	00BD	half
$\frac{3}{4}$	0190	00BE	driekwart
≡	240	2261	exact gelijk aan
≈	247	2248	ongeveer gelijk aan
≤	242	2264	groter dan of gelijk aan
≥	243	2265	kleiner dan of gelijk aan
√	251	221A	wortel
ⁿ	252	207F	tot de macht n
¹	0185	00B9	tot de macht 1
²	0178	00B2	tot de macht 2
³	0179	00B3	tot de macht 3
π	227	03C0	pi
°	248	00B0	graden
∞	236	221E	oneindig
μ	230	00B5	mu (micro-)
Σ	228	2211	som
©	0169	00A9	copyright

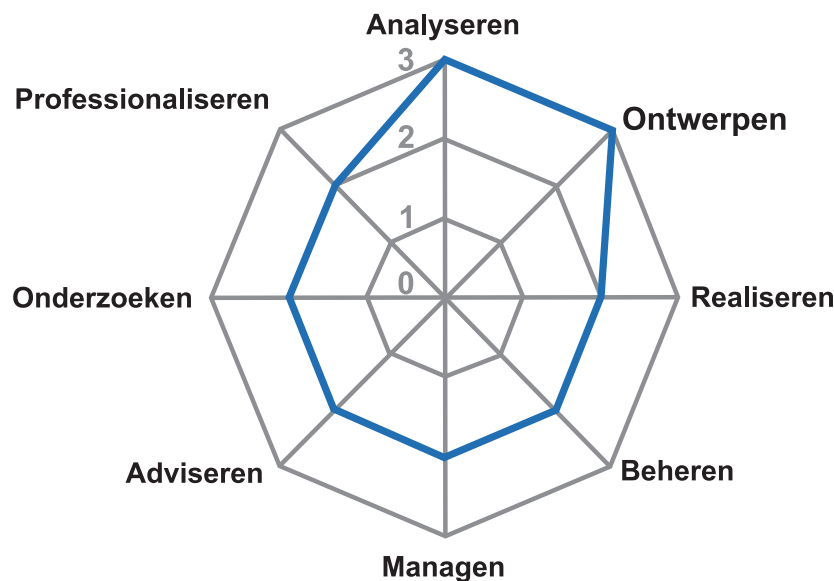
B. Bekende afgeleiden

$f(x)$	$\frac{df(x)}{dx}$
$\frac{1}{x}$	$-\frac{1}{x^2}$
e^x	e^x
$\ln(x)$	$\frac{1}{x}$
$\log(x)$	$\frac{1}{\ln(10) x}$
x^n	nx^{n-1}
10^x	$10^x \ln(10)$
$\sin(x)$	$\cos(x)$
$\cos(x)$	$-\sin(x)$
$\tan(x)$	$1 + \tan^2(x)$

C. Competenties bij het meten

Bij een hbo-studie Mechatronica horen competenties, die weer te geven zijn in een *spindiagram*. De cijfers daarin geven het niveau aan waarop een afgestudeerd mechatronicus minimaal eindigt.

SPIN Mechatronica



Op het terrein van *Meten & Weten* komen de genoemde competenties alle acht voor.

Analyseren Misschien wel de belangrijkste, omdat voor een meting een analyse van de onzekerheid nodig is. Onzekerheden zijn zelden vanzelfsprekend.

Ontwerpen Voordat je een meting uitvoert, moet je nadenken over hoe je dat gaat doen. Het kan de simpele keuze tussen een rolmaat en een duimstok zijn. Maar het bepalen van de valversnelling op een bepaalde plaats brengt heel wat meer keuzes met zich mee. Je ontwerpt een meetopstelling én een meetmethode.

Realiseren Iedereen die iets maakt of doet, realiseert iets. De uitvoering van een meting is een typisch geval van een realisatie. Wie dat deel niet goed voor elkaar heeft, komt tot onjuiste resultaten.

Beheren Een competentie die belangrijker en moeilijker is dan hij lijkt. Natuurlijk leg je de metingen ergens vast. Met pen en papier op een blaadje, of door in te tikken in een spreadsheet of ander programma. Hoe nummer je

de gegevens en hoe voorzie je ze van metadata (datum, tijd, door wie uitgevoerd) zodat je alles weet wat je later weten wilt? Hoe zorg je voor *backups* en voldoende toegankelijkheid? Hoe zorg je voor eventuele geheimhouding en de zekerheid dat data niet per ongeluk of expres veranderd worden?

Managen De uitvoering van een meting vraagt soms om *planning*. Wanneer heb je de data nodig en hoe ‘goed’ moet de meting zijn? Wanneer zijn de personen en de ruimte(s) beschikbaar waarin de metingen moeten worden uitgevoerd? En hoe zit dat met eventuele meetapparatuur?

Adviseren Misschien wel de meest algemene competentie, want adviseren kun je over bijna alles. Je bent pas competent op dit terrein als je voldoende kennis van zaken hebt. Maar dat alleen is niet voldoende. Het goed overbrengen op de doelgroep is net zo belangrijk.

Onderzoeken Voordat je een meetmethode ontwerpt, is er soms kennis nodig die je nog niet hebt. Bij een nieuwe of meer complexe meting kan dat het geval zijn. Het verwerven van die nieuwe kennis via (literatuur)onderzoek kan dan nodig zijn.

Professionaliseren Wie metingen verricht heeft, zal daar doorgaans over moeten communiceren. Hoe presenteer je de uitkomst zo dat de essentie van wat je te zeggen hebt goed overkomt? Dat een meting al dan niet strijdig is met iets anders, moet je expliciet vermelden. Wat de gevolgen daarvan zijn, moet je ook duidelijk maken. Veel metingen zullen aanleiding zijn tot reflecteren en leren. Volgende keer nóg beter.

D. Doorwerking van fouten

De basis: optellen, aftrekken, vermenigvuldigen, delen

functie	onzekerheid (afhankelijk)	onzekerheid (onafhankelijk)
$Z = A + B$	$\delta Z = \delta A + \delta B$	$\delta Z = \sqrt{(\delta A)^2 + (\delta B)^2}$
$Z = A - B$	$\delta Z = \delta A + \delta B$	$\delta Z = \sqrt{(\delta A)^2 + (\delta B)^2}$
$Z = A \times B$	$\frac{\delta Z}{ Z } = \frac{\delta A}{ A } + \frac{\delta B}{ B }$	$\frac{\delta Z}{Z} = \sqrt{\left(\frac{\delta A}{A}\right)^2 + \left(\frac{\delta B}{B}\right)^2}$
$Z = \frac{A}{B}$	$\frac{\delta Z}{ Z } = \frac{\delta A}{ A } + \frac{\delta B}{ B }$	$\frac{\delta Z}{Z} = \sqrt{\left(\frac{\delta A}{A}\right)^2 + \left(\frac{\delta B}{B}\right)^2}$

De formules voor aftrekken zijn hetzelfde als die voor optellen.

De formules voor delen zijn hetzelfde als die voor vermenigvuldigen.

Afgeleid: machten en factoren

functie	onzekerheid
$Z = A^n$	$\frac{\delta Z}{ Z } = \left n \frac{\delta A}{A} \right $
$Z = kA$	$\delta Z = k \delta A$

Machtsverheffen is vermenigvuldigen met het getal zelf.

Vermenigvuldigen met een exact getal is vermenigvuldigen met iets zonder onzekerheid.

Voor beide formules geldt dat er geen afhankelijke en onafhankelijke variant bestaat. Er is immers geen grootte om van afhankelijk te zijn.

Combinaties

functie	onzekerheid (afhankelijk)	onzekerheid (onafhankelijk)
$Z = k \frac{A}{B}$	$\frac{\delta Z}{ Z } = \frac{\delta A}{ A } + \frac{\delta B}{ B }$	$\frac{\delta Z}{Z} = \sqrt{\left(\frac{\delta A}{A}\right)^2 + \left(\frac{\delta B}{B}\right)^2}$
$Z = k \frac{A^n}{B^m}$	$\frac{\delta Z}{ Z } = \left n \frac{\delta A}{A}\right + \left m \frac{\delta B}{B}\right $	$\frac{\delta Z}{Z} = \sqrt{\left(n \frac{\delta A}{A}\right)^2 + \left(m \frac{\delta B}{B}\right)^2}$

Combinaties (vervolg)

functie	onzekerheid (onafhankelijk)
$Z = A + B - (C - D)$	$\delta Z = \sqrt{(\delta A)^2 + (\delta B)^2 + (\delta C)^2 + (\delta D)^2}$

Combinaties (vervolg)

functie	onzekerheid (onafhankelijk)
$Z = \frac{A \times B}{C \times D}$	$\frac{\delta Z}{Z} = \sqrt{\left(\frac{\delta A}{A}\right)^2 + \left(\frac{\delta B}{B}\right)^2 + \left(\frac{\delta C}{C}\right)^2 + \left(\frac{\delta D}{D}\right)^2}$
$Z = \frac{A^n \times B^m}{C^p \times D^q}$	$\frac{\delta Z}{Z} = \sqrt{\left(n \frac{\delta A}{A}\right)^2 + \left(m \frac{\delta B}{B}\right)^2 + \left(p \frac{\delta C}{C}\right)^2 + \left(q \frac{\delta D}{D}\right)^2}$

Omdat de formules lang worden, is alleen de kolom met de formules voor onafhankelijkheid weergegeven.

Wiskundige functies

functie f(x)	afgeleide	onzekerheid
$\frac{1}{x}$	$-\frac{1}{x^2}$	$\frac{\delta x}{x^2}$
e^x	e^x	$e^x \delta x$
$\ln x$	$\frac{1}{x}$	$\frac{\delta x}{x}$
$\log x$	$\frac{1}{\ln(10) x}$	$\frac{\delta x}{x}$
x^n	nx^{n-1}	$ nx^{n-1} \delta x$
10^x	$10^x \ln(10)$	$10^x \ln(10) \delta x$
$\sin x$	$\cos x$	$ \cos x \delta x$
$\cos x$	$-\sin x$	$ \sin x \delta x$
$\tan x$	$1 + \tan^2 x = \frac{1}{\cos^2 x}$	$\frac{\delta x}{\cos^2 x}$

Kettingregel

Als $f(x) = g(h(x))$ $f'(x) = g'(h(x)) \cdot h'(x)$

E. Enige SI-voorvoegsels

macht	voor	letter	naam	decimaal getal
10^{24}	yotta	Y	quadriljoen	1000000000000000000000000
10^{21}	zetta	Z	triljard	100000000000000000000000
10^{18}	exa	E	triljoen	10000000000000000000000
10^{15}	peta	P	biljard	1000000000000000000000
10^{12}	tera	T	biljoen	100000000000000000000
10^9	giga	G	miljard	10000000000000000000
10^6	mega	M	miljoen	1000000000000000000
10^3	kilo	k	duizend	1000
10^2	hecto	h	honderd	100
10^1	deca	da	tien	10

macht	voor	letter	naam	decimaal getal
10^{-24}	yocto	y	een quadriljoenste	0,000000000000000000000001
10^{-21}	zepto	z	een triljardste	0,000000000000000000000001
10^{-18}	atto	a	een triljoenste	0,000000000000000000000001
10^{-15}	femto	f	een biljardste	0,000000000000000000000001
10^{-12}	pico	p	een biljoenste	0,000000000000000000000001
10^{-9}	nano	n	een miljardste	0,0000000001
10^{-6}	micro	μ	een miljoenste	0,000001
10^{-3}	milli	m	een duizendste	0,001
10^{-2}	centi	c	een honderdste	0,01
10^{-1}	deci	d	een tiende	0,1

Deze bijlage heeft als titel “Enige SI-voorvoegsels”. Dat suggereert dat er meer zijn, maar dat is niet zo. Dit zijn de enige. Maar ik vond het leuk als elke bijlage een titel had met de beginletter van A, B, C, D, E, F, G en H.

F. Foutenvinders

De volgende personen hebben me blij kunnen maken door bij het maken van dit boek een of meer fouten te melden.

Tijdens het schrijven van de eerste uitgave in 2023:

Yannick van den Arend, Tommy Baas, Kaj van Beest, Tom Bos, Lars Hartman, Ilse van den Heuvel, Harm Jan van Holst, Guus Jansma, Bram Nijhoff, Daan Sijnja, Jimmy Scheltema, Koos Schoo, Jasper Waaijer, Michael van Weelde.

In de eerste gedrukte uitgave van 2023:

Bas van den Heuvel, Matthias Paul, Mark Schrauwen, Gerard Tuk.

G.Gebruikte boeken

- Barford, N. C. (1985). *Experimental Measurements: Precision, Error, and Truth* (second edition). John Wiley & Sons.
- Bell, Stephanie (1999) *A Beginner's Guide to Uncertainty of Measurement*. National Physical Laboratory
- Berendsen, H. J. C. (1997). *Goed meten met fouten: een introductie in de verwerking en presentatie van meetresultaten en de discussie van meetfouten*. Bibliotheek der RU Groningen
- Berendsen, H. J. C. (2011). *A Student's Guide to Data and Error Analysis (Student's Guides)* (first edition). Cambridge University Press.
- Bevington, P. & Robinson, K. D. (2002). *Data Reduction and Error Analysis for the Physical Sciences* (third edition). McGraw-Hill.
- Buonaccorsi, J. P. (2010). *Measurement Error: Models, Methods, and Applications (Chapman & Hall/CRC Interdisciplinary Statistics)* (first edition). Chapman and Hall/CRC.
- Elling, R., Andeweg, B., Baars, S., De Jong, J., & Swankhuisen, C. (2023). *Rapportagetechniek: schrijven voor lezers met weinig tijd*. Noordhoff.
- Fornasini, P. (2008). *The Uncertainty in Physical Measurements: An Introduction to Data Analysis in the Physics Laboratory* (2008 edition). Springer.
- Grabe, M. (2014). *Measurement Uncertainties in Science and Technology* (second edition). Springer.
- Hughes, I. & Hase, T. (2010). *Measurements and their Uncertainties: A practical guide to modern error analysis*. Oxford University Press.
- ISO (2008), *Guide to the Expression of Uncertainty in Measurement*, Geneva, International Organisation for Standardisation (JCGM 100:2008 GUM 1995 with minor corrections).
- Kahneman, D. (2012). *Thinking, fast and slow*. Penguin UK.
- Kaner, C., Bach, J., & Pettichord, B. (2001). *Lessons learned in software testing: A Context-Driven Approach*. Wiley.
- Kimothi, S. K. (2001). *Uncertainty of Measurements: Physical and Chemical Metrology: Impact & Analysis* (first edition). ASQ Quality Press.

- Kirkup, L. & Frenkel, R. B. (2006). *An Introduction to Uncertainty in Measurement: Using the GUM (Guide to the Expression of Uncertainty in Measurement)* (first edition). Cambridge University Press.
- Lyons, L. (1991). *A Practical Guide to Data Analysis for Physical Science Students* (first edition). Cambridge University Press.
- Merrin, J. (2017). *Introduction to Error Analysis: The Science of Measurements, Uncertainties, and Data Analysis* (first edition). CreateSpace Independent Publishing Platform.
- Morris, A. S. & Langari, R. (2020). *Measurement and Instrumentation: Theory and Application* (third edition). Academic Press.
- Peralta, M. (2012). *Propagation Of Errors: How To Mathematically Predict Measurement Errors*. CreateSpace Independent Publishing Platform.
- Ratcliffe, C. & Ratcliffe, B. (2016). *Doubt-Free Uncertainty in Measurement: An Introduction for Engineers and Students* (Softcover reprint of the original first edition 2015). Springer.
- Renkema, J., (2020). *Schrijfwijzer*. Boom.
- Salkind, N. J. (2016). *Statistics for People Who (Think They) Hate Statistics: Using Microsoft Excel 2016* (fourth edition). SAGE Publications.
- Sinek, S. (2009). *Start with Why: How Great Leaders Inspire Everyone to Take Action*. Penguin.
- Squires, G. L. (2001). *Practical Physics* (fourth edition). Cambridge University Press.
- Taylor, R. J. (1997). *An Introduction to Error Analysis: The Study of Uncertainties in Physical Measurements*. John R. Taylor.

De laatste pagina is leeg.

Bijna dan...