

HAAL MEER UIT JE HERSENEN

12 EN 19 NOVEMBER 2015

FETSJE BIJMA

VRIJE UNIVERSITEIT AMSTERDAM

f.bijma@vu.nl

I. INLEIDING: HET BREIN

De menselijke hersenen bestaan uit ongeveer 10^{13} zenuwcellen, de *neuronen*. Samen zorgen al deze neuronen voor het aansturen van het lichaam; bewegen, denken, voelen, zingen, etc. Dat gebeurt door middel van kleine elektrische boodschapjes tussen neuronen, de *actiepotentialen*. Ieder neuron heeft talloze wortels naar alle kanten. Via deze wortels krijgt hij berichten binnen, en via zijn *axon* (een belangrijke uitloper) verzendt hij zelf berichten.

Een zo'n elektrisch berichtje kan niet gemakkelijk gemeten worden. Maar wanneer veel neuronen tegelijkertijd berichten versturen in dezelfde richting, ontstaat een elektrische stroom die we aan de buitenkant van het hoofd kunnen meten. Met *electro-encephalografie* (EEG) kun je op de hoofdhuid de elektrische potentiaal meten. Deze potentiaal vertoont kleine schommelingen die veroorzaakt worden door de neurale activiteit. Een EEG scan wordt meestal op ca. 70 sensoren op de hoofdhuid gemeten. Op elk van deze sensors wordt de potentiaal bijvoorbeeld 1000 keer per seconde gemeten.

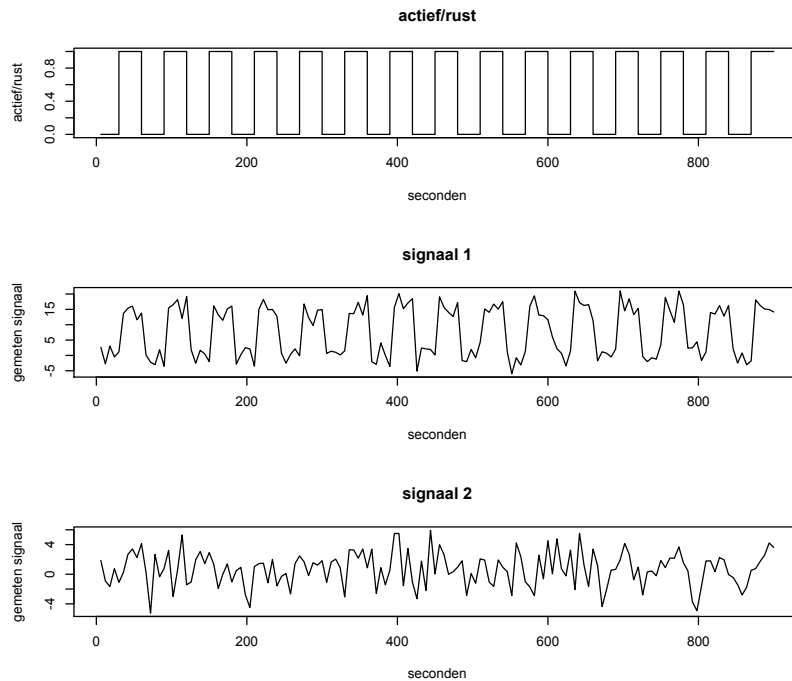
Een andere manier om hersenactiviteit in kaart te brengen is met behulp van een *Magnetic Resonance Imaging* (MRI) scanner. Bij het doorgeven van de berichten gebruiken de neuronen chemische stoffjes, de *neurotransmitters*. Nadat een neuron een berichtje doorgegeven heeft, moet dit stofje weer aangevuld worden. Daarvoor wordt er tijdelijk meer bloed toegevoerd naar de plek waar de overdracht van de boodschap heeft plaatsgevonden. Deze bloedtoevoer kan worden gemeten met een MRI-scanner. Een MRI-scanner maakt een 3D beeld: de hoeveelheid bloed wordt gemeten in 40.000 *voxels* in het brein. Een voxel is een kubusje met ribbe 3 mm. Een MRI-scanner is veel trager dan een EEG scanner: slechts een keer per 3 seconden wordt er een beeld gemaakt van het hoofd.

OPGAVE 1 Stel je een neuron in het brein even voor als een knikker van 1 cm doorsnee. Per neuron nemen we een knikker, en we leggen die in een lange streng. Hoelang wordt deze streng knikkers? Hoe vaak past deze streng rondom de aarde? (De omtrek van de aarde is ongeveer 40.000 km.)

OPGAVE 2 Bereken de inhoud van de hersenen (in liters) als je uitgaat van 40.000 voxels met ribbe 3 mm. Klopt dat ongeveer met jouw hoofd?

II. WELKE NEURONEN DOEN MEE?

In een gebruikelijk MRI-experiment worden actieve periodes afgewisseld met rustperiodes. Bijvoorbeeld afwisselend 30 seconden een knijpbeweging maken met je linkerhand en 30 seconden ‘niets doen’. Deze twee toestanden worden gedurende 15 minuten afgewisseld. Het idee is nu om te kijken welke voxels een soortgelijk aan/uit patroon vertonen. Als we in een bepaald voxel een signaal meten dat meer bloedtoevoer waarneemt tijdens de actieve periodes en juist minder tijdens de rustperiodes, dan mogen we concluderen dat die lokatie in de hersenen iets met de beweging te maken hebben. Neuronen op die plek zijn actiever wanneer we de beweging maken, ofwel dat zijn mogelijk de neuron die we gebruiken om de beweging aan te sturen. In figuur 1 zie je een voorbeeld van een activatiepatroon (de aan/uit functie voor de beweging) en twee voorbeeld signalen. Signaal 1 heeft duidelijk een relatie met het aan/uit patroon, terwijl dat bij signaal 2 niet direct duidelijk is.



Figuur 1: *Boven: de blokfunctie die de afwisseling tussen rust (0) en beweging (1) aangeeft. Midden: een voorbeeld van een signaal dat een relatie heeft met het activatiepatroon. Onder: een voorbeeld van een signaal dat niet een heel duidelijke relatie heeft met het activatiepatroon.*

In de volgende paragrafen gaan we naar wiskundige (of: statistische) methoden kijken waarmee we de relatie tussen signalen kunnen *kwantificeren*.

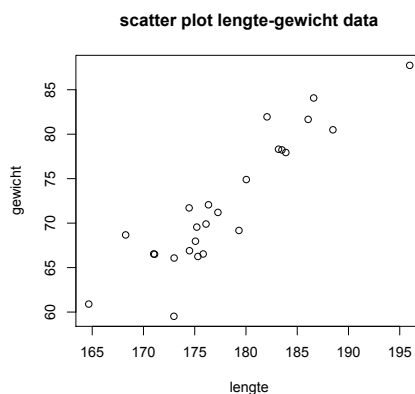
III. PUNTENWOLK

Stel we hebben de volgende metingen: $(x_1, y_1), \dots, (x_n, y_n)$. In de context van MRI is n het aantal tijdstippen waarop je meet, x_i de waarde van het blokpatroon en y_i de waarde van het signaal op tijdstip i . In het algemeen stelt het paar (x_i, y_i) een tweetal metingen voor aan het i^e “object”. Dat object kan een echt object zijn, maar het kan ook een tijdstip zijn, zoals bij MRI. In totaal zijn er n objecten die gemeten worden. Bijvoorbeeld: x_i is de lengte van de i^e persoon en y_i is het gewicht van de i^e persoon, voor $n = 25$ verschillende personen. Of: x_i is de gemiddelde temperatuur in december in de i^e stad en y_i is de hoogte boven de zeespiegel van de i^e stad, voor $n = 49$ verschillende steden wereldwijd. Of: x_i is de aanschafprijs van de i^e smartphone en y_i is de levensduur van de batterij van de i^e smartphone, voor $n = 14$ verschillende smartphones.

Deze gegevens kunnen we weergeven in een puntenwolk (ook wel: scatter plot, correlatiediagram, spreidingsdiagram, puntendiagram). We tekenen iedere waarneming (x_i, y_i) in het xy -vlak. Als voorbeeld kijken we naar de lengte-gewicht data van 25 personen in tabel 1. De puntenwolk staat getekend in figuur 2.

l	g	l	g	l	g	l	g	l	g
171	67	175	66	182	82	187	84	175	70
184	78	173	66	184	78	180	75	176	67
186	82	183	78	196	88	175	68	176	70
188	80	177	71	176	72	173	60	174	67
179	69	168	69	165	61	171	67	174	72

Tabel 1: *Lengte (l) in cm en gewicht (g) in kg van 25 personen.*



Figuur 2: *Correlatiediagram van de lengte-gewicht data: de punten (x_i, y_i) uit tabel 1 zijn getekend als punten in het xy -vlak.*

De punten in figuur 2 liggen rondom een (denkbeeldige) rechte lijn met positieve helling. Deze lijn kunnen we schrijven als

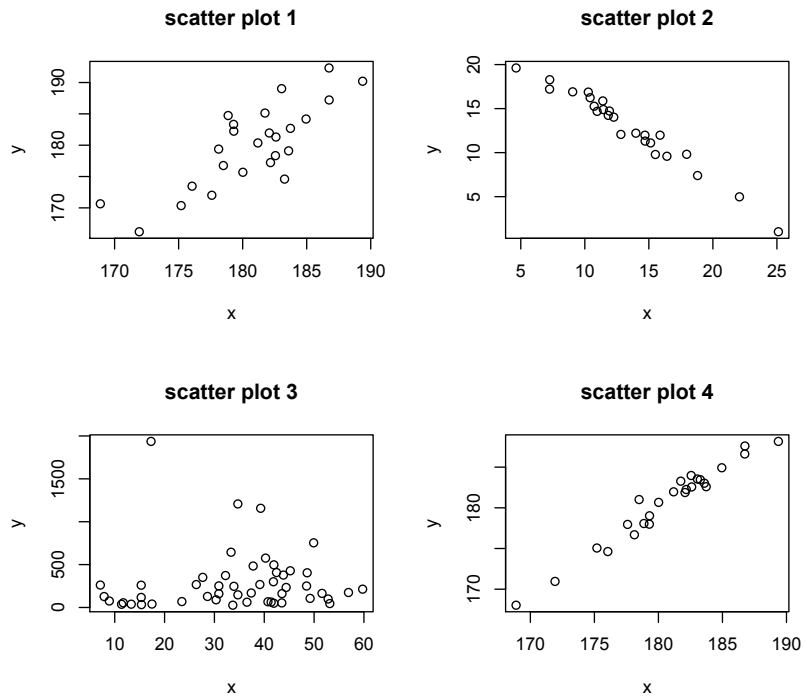
$$g = a + bl, \quad b > 0.$$

Hieruit kunnen we concluderen dat langere mensen gemiddeld zwaarder zijn dan kortere mensen. Alsof we dat nog niet wisten...

OPGAVE 3 Teken een goed passende rechte lijn in figuur 2. Wat zijn de waarden voor a en b in jouw lijn? Lig jouw eigen lengte-gewicht combinatie ongeveer op deze lijn?

OPGAVE 4 Bekijk de scatter plots in figuur 3.

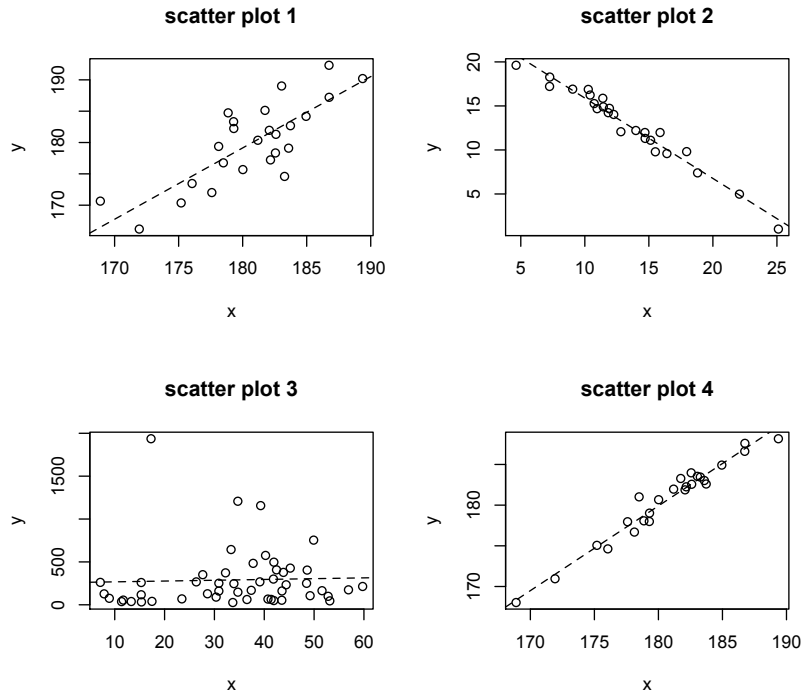
- a. In welke figuren liggen de punten goed rond een denkbeeldige rechte lijn? Is de helling daar positief of negatief?
 - b. Geef van de volgende vier beschrijvingen aan bij welke figuur de beschrijving zou kunnen horen.
 - i. x : gemeten lengte (cm) van oudste persoon binnen een eeneiige tweeling, y : gemeten lengte (cm) van jongste persoon binnen een eeneiige tweeling.
 - ii. x : gemeten lengte (cm) van oudste persoon binnen een twee-eiige tweeling, y : gemeten lengte (cm) van jongste persoon binnen een twee-eiige tweeling.
 - iii. x : gemeten temperatuur ($^{\circ}C$), y : gemeten tijd (min) nodig om 1 dl ijs te smelten bij gegeven temperatuur
 - iv. x : jaarlijkse regenval ($inch/j$) in VS staat, y : Gross Domestic Product (vergelijkbaar met Bruto Nationaal Product) van de staat ($10^9\$$)
-



Figuur 3: Vier scatter plots van onbekende data.

IV. CORRELATIE

De mate waarin de punten in een scatter plot variëren rond de denkbeeldige lijn kan uitgedrukt worden in de *correlatie* tussen x en y . In figuur 4 zijn dezelfde scatter plots te zien als in figuur 3, met daarin de best passende rechte lijn. (Met “best passende lijn” wordt de lineaire regressielijn bedoeld — dat voert te ver voor deze masterclass, maar leer je wel in de studie wiskunde!) Te zien is dat de variatie rond de lijn in scatter plot 1 groter is dan die in scatter plots 2 en 4. De lijn in scatter plot 3 lijkt niet zo zinvol, omdat er teveel variatie is rondom de lijn.



Figuur 4: De vier correlatiediagrammen uit figuur 3, nu met best passende lijn.

De formule voor de (steekproef) correlatiecoëfficiënt $r(x, y)$ is

$$r(x, y) = c \times \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})$$

De factor c is positief en werkt als een ijkingsfactor: deze factor zorgt ervoor dat $r(x, y)$ tussen -1 en +1 zit. De uitdrukking na het $\times$ teken bevat een *sommatieteken* ($\sum$). Deze uitdrukking staat voor:

$$\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) = (x_1 - \bar{x})(y_1 - \bar{y}) + (x_2 - \bar{x})(y_2 - \bar{y}) + \dots + (x_n - \bar{x})(y_n - \bar{y})$$

Het streepje in $\bar{x}$ en $\bar{y}$ betekent “gemiddelde nemen”. Dus

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_n}{n} \quad \text{en} \quad \bar{y} = \frac{y_1 + y_2 + \dots + y_n}{n}.$$

In de $r(x, y)$ wordt voor iedere meting (x_i, y_i) bepaald hoeveel de x_i van de gemiddelde x -waarde ($\bar{x}$) afligt, en net zo voor de y -waarde. Er zijn vier mogelijkheden:

- i. als $x_i \geq \bar{x}$ en $y_i \geq \bar{y}$, dan $(x_i - \bar{x})(y_i - \bar{y}) \geq 0$
- ii. als $x_i \geq \bar{x}$ en $y_i \leq \bar{y}$, dan $(x_i - \bar{x})(y_i - \bar{y}) \leq 0$
- iii. als $x_i \leq \bar{x}$ en $y_i \geq \bar{y}$, dan $(x_i - \bar{x})(y_i - \bar{y}) \leq 0$
- iv. als $x_i \leq \bar{x}$ en $y_i \leq \bar{y}$, dan $(x_i - \bar{x})(y_i - \bar{y}) \geq 0$.

Als de punten (x_i, y_i) precies op een rechte lijn met positieve helling liggen, dan komen alleen geval (i.) en (iv.) voor. In dat geval is $r(x, y) = 1$. Als de punten precies op een rechte lijn met negatieve helling liggen, dan komen alleen geval (ii.) en (iii.) voor. In dat geval is $r(x, y) = -1$. In alle andere gevallen krijgen we combinaties van de gevallen (i.) t/m (iv.) en ligt $r(x, y)$ tussen -1 en +1.

Voor de scatter plots in figuur 3 ga je nu de correlatiecoëfficiënten berekenen. Voor het maken van deze opgaven heb je de software Rstudio nodig. Deze software is al geïnstalleerd op de computers in de computerzaal voor de masterclass. Wil je er thuis ook mee werken? Je kunt het gratis downloaden. Installeer eersr *R* van de site www.r-project.org en vervolgens Rstudio van de site www.rstudio.com.

OPGAVE 5 Bepaal de uitkomst van de volgende sommaties:

- i. $\sum_{x=1}^4 x$
- ii. $\sum_{i=1}^5 i$
- iii. $\sum_{k=1}^6 k$
- iv. $\sum_{r=1}^7 r$

Kun je een algemene formule bedenken voor $\sum_{x=1}^n x$?

OPGAVE 6 In deze opgave controleer je de uitkomsten van opgave 5 met Rstudio. Start Rstudio op. In het window **linksboven** typ je je *R*-commando's. Het venster **rechtsboven** toont alle waarden in het geheugen. Het venster **links-onder** laat zien welke commando's gebruikt zijn en **rechtsonder** verschijnen de grafieken.

- i. Bereken $\sum_{x=1}^4 x$ door linksboven `sum(1:4)` te typen. Een getypt commando kun je *runnen* door Ctrl+Enter te drukken. De uitkomst verschijnt linksonder.
- ii. Bereken net zo $\sum_{i=1}^5 i$, $\sum_{k=1}^6 k$ en $\sum_{r=1}^7 r$. Kloppen je uitkomsten in opgave 5?
- iii. Voor het vinden van een algemene formule maken we een grafiek van de punten $(x, \sum_{x=1}^n x)$ voor $x = 1, 2, 3, \dots, 100$. Typ en *run* de volgende R-code:

```
x=1:100
y=cumsum(x)
x=plot(x,y)
```

Druk op het knopje Zoom boven de grafiek om een groot grafiekvenster te zien. Welk functievoorschrift past hierbij? Een rechte lijn $y = ax + b$? Een parabool $y = ax^2 + bx + c$? Een derdegraadspolynoom $y = ax^3 + bx^2 + cx + d$? Vind ook de waarden van a , b , etc.

OPGAVE 7 Op www.few.vu.nl/~fetsje/masterclass vind je de R-commando's die je kunt gebruiken voor deze masterclass. Kopieer alle tekst naar het venster linksboven. Deze R-code kun je nu *runnen* met Ctrl+Enter. Dat kan per regel, maar je kunt ook meerdere regels selecteren en dan Ctrl+Enter drukken om alle regels tegelijk te runnen. Regels die beginnen met een “#” zijn geen commando, maar bevatten informatie. Je kunt deze regels veilig *runnen*, want R slaat ze automatisch over vanwege het #-teken.

- i. Run de code bij II. WELKE NEURONEN DOEN MEE? en maak figuur 1 uit de syllabus na.
- ii. Run de code bij III. PUNTENWOLK en bereken de correlatiecoëfficiënt van de lengte-gewicht data, zoals dat staat in de code onder IV. CORRELATIE.
- iii. Bereken vervolgens op soortgelijke wijze de correlatiecoëfficiënt voor elk van de vier scatter plots in figuur 3. Welke waarde(n) geven aan dat er een goed lineair verband is? Is dat positief of negatief? Welke waarde(n) geven aan dat er weinig lineair verband is? Kloppen de uitkomsten met wat je had verwacht?

Als het goed is geven de waarden die je in opgave 7 hebt berekend precies aan wat we ook al hadden gezien: in scatter plot 2 en 4 liggen de punten het best op een rechte lijn ($r(x, y)$ heel dicht bij ± 1). In scatter plot 1 is ook aardig sprake van een rechte lijn ($r(x, y) = 0.81$) terwijl in scatter plot 3 de rechte lijn niet zinvol lijkt ($r(x, y) = 0.04$).

Waar ligt de grens tussen “duidelijk een rechte lijn” en “een rechte lijn niet zinvol”? Dat is lastig te zeggen op basis van alleen een correlatiediagram. Daarvoor is een statistische toets nodig, die in het vervolg van deze masterclass wordt behandeld.

V. STATISTISCHE TOETSEN: DE ONEERLIJKE DOBBELSTEEN

Met een statistische toets kan een bewering getoetst worden op basis van beschikbare data. Als voorbeeld kijken we naar een dobbelsteen. Stel we gooien 5 keer met de dobbelsteen en gooien alle 5 keren zes ogen. Is dat niet een beetje tè toevallig voor een zuivere dobbelsteen? Wat is de kans op deze uitkomst als de dobbelsteen eerlijk is? Met een zuivere dobbelsteen is de kans op zes ogen gelijk aan $1/6$. De kans dat we twee keer achter elkaar zes ogen gooien is dan $1/36$, etc. De kans op 5 keer achter elkaar zes ogen is $(\frac{1}{6})^5 = 0.0001$. Dat is inderdaad tè toevallig (maar niet onmogelijk!). Bij een statistische toets hanteert men een grenswaarde voor de kans, bijvoorbeeld van 1%. Als, onder aanname dat de dobbelsteen zuiver is, de kans op de waargenomen uitkomst (of extremer) kleiner is dan de grenswaarde, dan “verwerpen we de uitspraak” dat de dobbelsteen zuiver is. In dat geval nemen we aan dat de dobbelsteen oneerlijk is.

In het algemeen wordt een statistische toets gebruikt om een bewering (bijvoorbeeld “deze dobbelsteen is zuiver”) te toetsen. Deze bewering wordt de *nulhypothese* genoemd. We hebben nodig:

- waargenomen uitkomst (“5 keer zes ogen achter elkaar gooien”)
- de kans op de waargenomen uitkomst of extremer, onder de aanname dat de nulhypothese waar is (“0.0001”). Deze kans heet de *p-waarde*.
- een grenswaarde voor de *p-waarde* (“1%”).

De uitslag van een statistische toets wordt gebaseerd op de *p-waarde*. Er zijn twee mogelijkheden:

- **Sterke conclusie:** De *p-waarde* is *lager* dan de grenswaarde. Dan is de waarneming heel onwaarschijnlijk als de nulhypothese waar is. Daarom verwerpen we in dat geval de nulhypothese, en nemen we aan dat het tegenovergestelde waar is.
- **Zwakke conclusie:** De *p-waarde* is *groter* dan de grenswaarde. Dan is de waarneming niet heel uitzonderlijk als de nulhypothese waar is. In dat geval hebben we niet genoeg bewijs om de nulhypothese te verwerpen. Dan laten we de waarheid in het midden. (Dat betekent dat we in dit geval niet hebben “bewezen” dat de nulhypothese waar is! Dat kan namelijk niet met 100% zekerheid.)

OPGAVE 8 Stel we gooien 5 keer met een eerlijke dobbelsteen. Bereken de kans op 0, 1, 2, 3 en 4 maal zes ogen binnen deze vijf worpen. Hoe vaak moet men zes ogen gooien om een *p-waarde* lager dan 1% te krijgen?

aantal maal 6 ogen	frequentie	geschatte kans op uitkomst
0	4041	
1	4015	
2	1598	
3	313	
4	32	
5	1	
totaal	10.000	1

Tabel 2: *Frequentie van aantal keer zes ogen binnen 5 worpen bij 10.000 herhaalde pogingen.*

Om een statistische toets in de praktijk te kunnen gebruiken, is het nodig dat we de p -waarde, ofwel de kans op de waargenomen uitkomst, kunnen bepalen. In het geval van de dobbelsteen is dat vrij eenvoudig, maar in de praktijk vereist dit soms veel rekenwerk. Als dit rekenwerk te complex wordt, kan het vaak vervangen worden door *simulatie*. Simulatie houdt in dat je met de computer uitkomsten nabootst, onder aanname dat de nulhypothese waar is. In het voorbeeld van de dobbelsteen kun je de computer 5 keer een dobbelsteen laten gooien (denkbeeldig), en tellen hoeveel keer zes ogen is gegooid. De mogelijke uitkomsten zijn dan 0, 1, 2, 3, 4 of 5 keer zes ogen. Dit experiment herhaal je heel vaak, bijvoorbeeld 10.000 keer. We krijgen dan de resultaten in tabel 2.

OPGAVE 9 Bereken met behulp van de tweede kolom in tabel 2 een schatting voor de kans op 0, 1, ..., 5 keer zes ogen. Komen deze waarden goed overeen met je antwoord bij opgave 8?

Zoals je ziet kun je de kansverdeling met simulatie aardig benaderen. Simulaties geven echter geen unieke uitkomst. Als we de simulatie voor de dobbelsteen opnieuw doen, weer met 10.000 maal 5 keer gooien, dan vinden we bijvoorbeeld de uitkomsten in tabel 3. Dat de twee tabellen verschillend zijn, illustreert dat simulatie niet hetzelfde is als een precieze berekening. Met simuleren *benader* je een precieze berekening. In zo'n benadering wordt altijd een fout gemaakt. Deze fout wordt kleiner naarmate je meer waarnemingen simuleert.

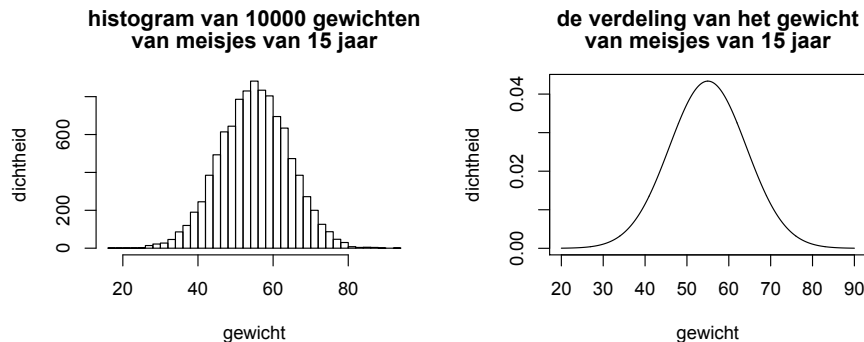
Veel gevallen in de praktijk verschillen in twee opzichten van deze dobbelsteensituatie. Ten eerste kan de exacte kans vaak niet berekend worden, omdat het te complex is. In zo'n geval kan met simulatie een p -waarde bepaald worden. In het geval van de dobbelsteen zou de p -waarde voor 5 keer zes ogen gelijk zijn aan 0.0001 (op basis van tabel 2) of 0.0003 (op basis van tabel 3). In beide gevallen is de p -waarde veel kleiner zijn dan de grenswaarde van 1% en trekken we de sterke conclusie: de dobbelsteen is niet zuiver. Het tweede verschil is dat er dikwijls meer dan 6 mogelijke uitkomsten zijn. Dan kunnen de uitkomsten niet op een overzichtelijke manier in een tabel worden gezet. Wanneer er veel meer uitkomsten mogelijk zijn, kunnen we de kansen beter weergeven in een plaatje (nl. een histogram).

aantal maal 6 ogen	frequentie
0	4070
1	4093
2	1502
3	292
4	40
5	3
totaal	10000

Tabel 3: *Frequenties als in tabel 2 in een tweede simulatie.*

VI. STATISTISCHE TOETSEN: GEWICHTEN VAN 15-JARIGE MEISJES

Het gewicht van meisjes van 15 jaar in Nederland is gemiddeld ongeveer 55 kg. Daarom heen zit fluctuatie, die redelijk beschreven kan worden met een normale verdeling, als in figuur 5.

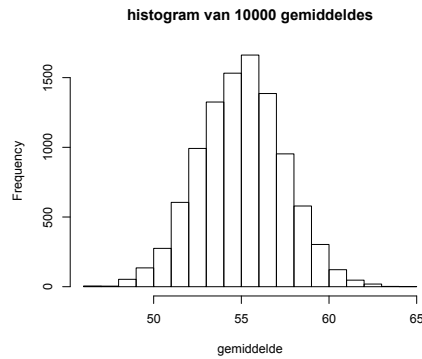


Figuur 5: *Links: histogram van de gewichten van 10.000 meisjes van 15 jaar. Op de y-as het aantal meisjes met een gewicht in het interval op de x-as van de staaf. Rechts: de normale verdeling die de variatie in de gewichten beschrijft.*

Links in figuur 5 is het histogram getekend van de gewichten van 10.000 meisjes van 15 jaar. De x -as is in intervallen van lengte 2 verdeeld, en per interval is het aantal gewichten in dat interval geteld. De aantallen staan op de y -as uitgezet. Het is gebruikelijk om de y -as met een factor te vermenigvuldigen, zodat de oppervlakte van alle staafjes samen gelijk is aan 1. Als we dat doen en we meten steeds meer meisjes, dan wordt het histogram steeds nauwkeuriger. Als we oneindig veel meisjes zouden kunnen meten, en intervallen steeds kleiner kiezen, dan krijgen we uiteindelijk de normale dichtheid als in de rechter figuur in figuur 5. Dit is de normale verdeling met gemiddelde 55 en standaarddeviatie 9.2.

Nu gaan we met een statistische toets onderzoeken of meisjes uit een heel sportieve klas (met maar liefst 8 uur gym per week) gemiddeld lichter zijn dan “gemiddelde meisjes” in Nederland. We meten de gewichten van 14 meisjes van 15 jaar: 39.4, 40.8, 42.2, 46.3, 46.4, 47.0, 47.9, 51.5, 51.5, 52.8, 56.4, 60.5, 65.9 en 69.4 kg. We toetsen de nulhypothese: “Het gewicht van deze sportieve meisjes is normaal verdeeld met gemiddelde 55 en standaarddeviatie 9.2”. Om nu een p -waarde te bepalen, kunnen we kijken naar het waargenomen gemiddelde. Dat is 51.3 kg. Is dat te laag om de bewering te geloven, of zou het wel kunnen? De p -waarde geeft het antwoord op die vraag. De p -waarde kunnen we bepalen met behulp van simulatie.

We trekken daartoe 10.000 keer een steekproef van 14 gewichten uit de normale verdeling met gemiddelde 55 en standaarddeviatie 9.2. Van iedere steekproef berekenen we het gemiddelde. Het histogram van deze 10.000 gemiddelden staat in figuur 6.



Figuur 6: *Het histogram van de gemiddelden van de 10.000 steekproeven ter grootte 14 uit de normale verdeling met gemiddelde 55 en standaarddeviatie 9.2.*

De p -waarde kunnen we nu schatten op basis van deze simulatieuitkomsten. We tellen hoeveel van de 10.000 gemiddelden lager zijn dan 51.3. Dat zijn er 662. Dat betekent dat, onder de aanname dat de nulhypothese waar is, de kans op het waargenomen gemiddelde van 51.3 *of lager* gelijk is aan $662/10000=0.0662$. Dus onze waarneming is niet héél toevallig, en we kunnen de nulhypothese op basis van deze meting niet verwerpen. We concluderen niets (de zwakke conclusie) in dit geval. Meisjes uit sportieve klassen zouden lichter kunnen zijn, maar in onze meetgegevens zit daarvoor niet voldoende “bewijs”.

VII. STATISTISCHE TOETSEN SAMENGEVAT

In de voorbeelden in de vorige twee paragrafen hebben we simulatie gebruikt om de p -waarde van een uitkomst te bepalen. Als uitkomst hebben we genomen: het aantal maal zes ogen (bij de dobbelsteen) en het gemiddelde gewicht van 14 meisjes (bij de gewichten). Zo'n samenvattende grootheid van de uitkomst heet een *toetsingsgrootheid*. We hebben niet naar alle 14 gewichten gekeken, maar slechts naar het gemiddelde. Dat bevat voldoende informatie. Voor een statistische toets is altijd een toetsingsgrootheid nodig, die een uitkomst beschrijft in slechts één getal.

Samengevat heeft een statistische toets de volgende stappen:

1. stel een nulhypothese op
2. bepaal een toetsingsgrootheid
3. bepaal de p -waarde (met berekening of met simulatie)
4. bepaal de conclusie van de toets.

OPGAVE 10 Neem aan dat het gewicht van meisjes van 15 jaar normaal verdeeld is met gemiddelde 55 en standaarddeviatie 9.2, zoals in figuur 5. Stel nu dat we het gewicht van 25 meisjes meten die juist weinig lichaamsbeweging krijgen. Het gemiddelde van deze 25 gewichten is 59.5 *kg*. Voer nu de statistische toets uit (zie de *R*-code voor een voorbeeld) om te bepalen of deze meisjes in gemiddeld gewicht afwijken van “gewone meisjes”.

VIII. TERUG NAAR DE HERSENEN

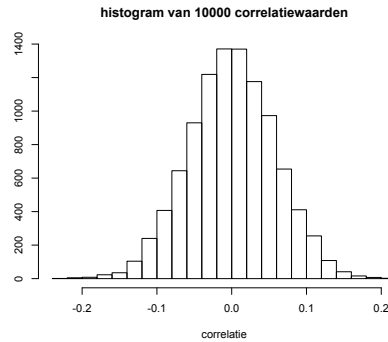
Om te bepalen welke gebieden in de hersenen te maken hebben met een bepaalde beweging willen we toetsen welke signalen significant correleren met het activatiepatroon. In andere woorden: we nemen als toetsingsgrootheid de correlatiecoëfficiënt tussen het gemeten signaal in een voxel en het blokpatroon dat de afwisseling tussen rust en beweging beschrijft (zie figuur 2).

Als de signalen onafhankelijk zijn van elkaar dan is de correlatie tussen x -waarden en y -waarden *ongeveer* gelijk aan 0. Dit komt overeen met een grote spreiding in de puntenwolk. Als de correlatie tussen x -waarden en y -waarden *significant* van 0 afwijkt dan is de spreiding rondom de denkbeeldige rechte lijn in een puntenwolk niet al te groot is. Wat is te groot, en wat niet? De grens daartussen bepalen we met een statistische toets.

Voor de statistische toets hebben we nodig:

- de nulhypothese: de correlatie wijkt niet significant af van 0.
- een toetsingsgrootheid: de correlatie $r(x, y)$ tussen de x -waarden en de y -waarden.
- een p -waarde (met behulp van simulatie)

Voor het bepalen van de p -waarden moeten we 10.000 correlatiewaarden berekenen op basis van onafhankelijke x en y -waarden. Daartoe simuleren we 10.000 keer een rij x -waarden en onafhankelijk daarvan een rij y -waarden. De correlatiecoëfficiënt $r(x, y)$ van deze gesimuleerde waarden varieert rondom nul. Een histogram van deze gesimuleerde $r(x, y)$ -waarden staat in figuur 7.



Figuur 7: *Het histogram van de gesimuleerde correlatiewaarden $r(x, y)$ op basis van onafhankelijke x - en y -waarden.*

Om te bepalen of het gemeten hersensignaal in een bepaald voxel significant (positief of negatief) correleert met het activatiepatroon bepalen we de correlatiecoëfficiënt voor het voxel en bepalen we een p -waarde uit de gesimuleerde $r(x, y)$ -waarden. Als de p -waarde kleiner is dan 1%, dan kunnen we concluderen dat het betreffende voxel, ofwel het betreffende hersengebied, te maken heeft

met de beweging. Als de correlatie positief is dan is er tijdens de beweging meer hersenactiviteit in dit gebied. Is de correlatie negatief, dan is er juist minder activiteit tijdens de beweging. Dat klinkt misschien gek. Toch gebeurt dit op veel plekken in de hersenen. Een aantal neuronen is juist actief als het lichaam in rust is, en andersom. Deze neuronen vertonen tijdens de beweging dus juist minder activiteit!

OPGAVE 11 Gebruik de *R*-code van VIII. TERUG NAAR DE HERSENEN om de gesimuleerde correlatiewaarden te maken met Rstudio. De data van de voxels staat op de website van de masterclass. Het activatiepatroon en de data van 20 voxels kun je het snelst kopiëren en plakken in het venster **linksonder** (dus niet linksboven). Bepaal voor welke voxels het signaal significant correleert met het activatiepatroon. Is de correlatie voor die voxels negatief of positief?

IX. NIEUWSGIERIG GEWORDEN?

Vond je deze masterclass leuk en ben je nieuwsgierig geworden naar de diepere wiskundige theorie hierachter? Dat (en nog veel meer!!) leer je allemaal in de bacheloropleiding Wiskunde. De wiskundige techniek die hier achter zit heet *lineaire regressie*. Dat wordt heel veel gebruikt in wetenschappelijke studies.