

MC SYLLABUS 34

J. HEMELRIJK

**ORIENTERENDE CURSUS
MATHEMATISCHE STATISTIEK**

MATHEMATISCH CENTRUM

AMSTERDAM 1977

AMS(MOS) subject classification scheme (1970): 62-01,62E15,62F05,62F10,62G05,62G10

ISBN 90 6196 139 4

VOORWOORD

De eerste uitgave van deze cursus, met rapportnummer S200, verscheen in 1956 in gestencilde vorm. Dat is geruime tijd geleden en de cursus is dan ook gedeeltelijk verouderd, zoals o.a. blijkt uit voorbeeld 3 op blz. van de inleiding: "De wachttijd van een auto bij de Hembrugpont". Wie bekommert zich daar nu nog om! Een ander punt, dat althans volgens de schrijver op veroudering wijst, is, dat gekozen is voor de axiomatische opzet van de kansrekening, zij het geruggesteund door de wet van de grote aantallen. Zonder de daarbij behorende frequentie-interpretatie te verloochenen komt het heden de schrijver beter voor de opzet van de elementaire kansrekening en statistiek te baseren op de definitie van Laplace, aangevuld met het gebruik van een aselector voor het nemen van aselechte steekproeven^{*)}.

De vraag naar S200 bleef echter, lang nadat de syllabus in verschillende drukken verscheen en uitverkocht was, aanhouden. Daarom werd ten slotte besloten de cursus vrijwel ongewijzigd te herdrukken, nu als deel van de MC Syllabus reeks. Hiertoe was nogal wat voorbereiding nodig in de vorm van correcties, kleine wijzigingen en aanvullingen. Dit werk werd verricht door A. Wolowitsj, die hiervoor de dank van de schrijver oogst.

Amsterdam, december 1976

J. Hemelrijk.

^{*)} J. HEMELRIJK, *Underlining random variables*, Statistica Neerlandica 20 (1966) 1-7,
 J. HEMELRIJK, *Back to the Laplace definition*, Statistica Neerlandica 22 (1968) 13-21.
 B. VAN ROOTSELAAR & J. HEMELRIJK, *Back to "Back to the Laplace definition"*, Statistica Neerlandica 23 (1969), 87-89.

<i>Voorwoord</i>	<i>i</i>
<i>Inhoud</i>	<i>iii</i>
<i>Literatuur</i>	4
Hoofdstuk I INLEIDING	1
II HET EXPERIMENTELE FUNDAMENT VAN DE STATISTIEK	10
III EIGENSCHAPPEN VAN FREQUENTIEQUOTIËNTEN.	14
IV HET KANSBEGRIIP.	24
V VOORBEELDEN VAN HET BEREKENEN VAN KANSEN.	36
VI HET METEN VAN KANSEN.	45
VII STOCHASTISCHE GROOTHEDEN, KANSVERDELINGEN	47
VIII ONAFHANKELIJKHEID VAN STOCHASTISCHE GROOTHEDEN. . . .	61
IX MATHEMATISCHE VERWACHTING EN MOMENTEN	69
X DE ONGELIJKHEID VAN BIENAYMÉ-TCHEBYCHEFF EN DE THEORETISCHE WET VAN DE GROTE AANTALLEN	85
XI FUNCTIES VAN STOCHASTISCHE GROOTHEDEN; CENTRALE LIMIETSTELLING	90
XII SCHATTINGSTHEORIE; MEEST AANNEMELIJKE SCHATTINGEN . .	104
XIII TOETSINGSTHEORIE; BINOMIALE TOETSEN	136
XIV ENIGE GANGBARE TOETSEN.	156
XV BETROUWBAARHEIDSGRENZEN	202
APPENDIX.	223
TABEL 1 t/m 5c.	227
REGISTER.	234

HOOFDSTUK I

INLEIDING

Gaandeweg is de wiskunde als hulpwetenschap doorgedrongen in velerlei gebieden van wetenschap en praktijk, gedeeltelijk is het ontstaan ervan zelfs direct aan bepaalde praktische problemen verbonden. Meetkunde, algebra, analyse, enz. worden als routinemiddelen gehanteerd in de natuurkunde, sterrekunde, scheepvaartkunde, en waar al niet meer. In de regel - en tot voor betrekkelijk kort vrijwel steeds - wordt de wiskunde gebruikt om problemen van deterministische aard op te lossen. De elementaire natuurkunde b.v. houdt zich bezig met verschijnselen, die zich onder gelijke omstandigheden steeds op dezelfde wijze voordoen. Voor toepassing van wiskundige theorieën bedient men zich dan van een *wiskundig model* van dat (kleine) deel van de waarneembare werkelijkheid, dat men onderzoekt. Zo is een rechte lijn b.v. de model-voorstelling van de draad van een schietlood, van een muur of van een schrikdraad; men werkt met getalwaarden voor het gewicht of de lengte van een voorwerp, de grootte van een kracht, de lading van een electron, alsof deze begrippen ondubbelzinnig bepaald zijn en een exact bepaalbare grootte bezitten.

Dergelijke modellen leiden dan tot exact gedetermineerde uitkomsten, hetgeen in de praktijk neerkomt op pertinente uitspraken en voorspellingen: "Op die dag zal van zo tot zo laat een zonsverduistering optreden", "de omtrek van de aard-eauator is 40.000 km", etc. Weliswaar bezitten deze uitspraken slechts een beperkte precisie - waar men zich ook wel van bewust is - maar zolang deze voldoende is voor het doel, waarop het onderzoek gericht is, zijn de uitspraken in hun deterministische vorm bruikbaar.

Vele verschijnselen echter laten zich niet - of althans niet in voldoende mate - in een dergelijk model vangen. Deterministische methoden leiden tot gedetailleerde voorspellingen en vele zaken laten zich niet - of althans nog niet - in detail voorspellen. Voorbeelden hiervan zijn:

1. Het weer van morgen (of, in sterkere mate, van volgende week).
2. Het aantal verkeersongelukken in Amsterdam gedurende de eerstvolgende Paasweek.
3. De wachttijd van een auto bij de Hembrugpont.
4. Het jaarlijkse nationale inkomen.
5. De werkzaamheid van een geneesmiddel.
6. Het cijfer voor meetkunde van een examinandus, enz.

Hoewel dit soort "onzekere" verschijnselen zich niet in detail laat voorspellen, tast men er toch niet geheel over in het duister. Indien men voorspelt, dat het morgen zal regenen, heeft men in Nederland vaak - zij het niet altijd - gelijk en nog vaker indien men voorspelt, dat het weer morgen "hetzelfde" zal zijn als vandaag. Indien men beschikt over de ongevalstatistiek van vorige jaren kan men een goede slag slaan naar het aantal ongelukken in de niet te verre toekomst. Onderzoekingen omtrent de werkzaamheid van een geneesmiddel verschaffen een globale indruk daarvan; maar of het in een bepaald geval zal helpen of niet kan men niet met zekerheid voorspellen.

De statistiek houdt zich bezig met dit soort onzekere verschijnselen. Het is de wetenschap van het "hoe vaak" als de vraag "wanneer" niet zou worden beantwoord. Zo kan men b.v. bij een reeks worpen met een dobbelsteen niet voorspellen bij welke worpen een 6 zal optreden, maar wel (bij goede benadering) hoe vaak dit in een lange reeks zal gebeuren. Onvoorspelbaarheid in detail impliceert nog niet onvoorspelbaarheid in het groot. Dit is het fundamentele (experimentele) feit, waarop de statistiek gebaseerd is.

Het toepassingsgebied van de wiskunde is door het ontwikkelen van wiskundige methoden voor de behandeling van onzekere verschijnselen zeer aanzienlijk uitgebreid. Om slechts enkele gebieden te noemen waar de statistiek een belangrijk hulpmiddel is gebleken: biologie, medicijnen, landbouw, industrie, geodesie, astronomie, verzekering, economie, sociologie, bevolkingsleer, strategie en kernphysica.

Het is duidelijk, dat men in detail onvoorspelbare verschijnselen niet in een deterministisch model kan vangen. De statistiek bedient zich dan ook vaak van modellen, waarin de onzekerheid een essentiële element is, n.l. van dat deel van de zuivere wiskunde, dat de *kansrekening* of

waarschijnlijkheidsrekening genoemd wordt en dat een onderdeel van de maattheorie is. Het essentiële praktische verschil met de "klassieke" beschouwingwijze is, dat men bij toepassing van de statistiek afziet van de eis, dat *iedere* uitspraak of voorspelling juist is - iets wat men bij "klassieke" methoden weliswaar evenmin bereikt, maar waarop deze methoden toch gericht zijn -, maar dat men tevreden is als dit in de regel, b.v. in 95% of 99% van de gevallen, waarin men statistiek toepast, het geval is. De onzekerheid van de menselijke kennis wordt in de statistiek erkend, opgenomen en gekwantificeerd.

Verdere uitwerking van deze grondgedachte vindt plaats in latere hoofdstukken. Ter illustratie geven wij hieronder enkele voorbeelden van problemen, die met behulp van statistische methoden onderzocht kunnen worden. Bij ieder van deze problemen zal, zoals eigenlijk vanzelf spreekt, de statistische analyse van het probleem, gebaseerd moeten worden op waarnemingen, die de numerieke gegevens voor de berekeningen moeten leveren.

1. Bij het ontwikkelen van een methode voor het doen van weersvoorspellingen is statistische analyse van grote hoeveelheden waarnemingsmateriaal een belangrijk hulpmiddel.
2. Het werk van de Delta-commissie ter bepaling van de hoogte van dijken, die nodig zijn om in de toekomst de kans op overstromingen zeer gering te maken, berust voor een deel op statistische analyse van de hoogwaterstanden, die in het verleden waargenomen zijn.
3. Voor het onderzoeken van de al of niet werkzaamheid van een geneesmiddel en bij het vergelijken van twee of meer geneesmiddelen zijn experimenten nodig, waarvan de uitkomsten statistisch geanalyseerd dienen te worden.
4. Hetzelfde geldt voor talrijke technische vraagstukken, b.v. het onderzoeken van de invloed, die verschillende factoren, zoals de soort en de hoeveelheid van de brandstof en snelheid en hoeveelheid van de toegevoegde lucht, hebben op de straling van een vlam in een vlamoven.
5. Landbouwkundig onderzoek van de invloed van factoren, zoals grondbewerking, bemesting en variëteit op de opbrengst van een akker berusten tegenwoordig vrijwel steeds op statistische methoden, waarbij vooral ook aandacht wordt besteed aan een doeltreffende opzet der - gewoonlijk zeer tijdrovende - experimenten.

Hoewel nog vele andere voorbeelden genoemd zouden kunnen worden - van de walvisvangst tot het breken van garen in de textielindustrie, van de maximum snelheid van bromfietsen tot het schatten van de ouderdom van aardlagen - volstaan wij hier met de bovenstaande.

Tot slot van deze inleiding merken wij nog slechts op, dat de eerste stap van de statistische verwerking van waarnemingsmateriaal gewoonlijk bestaat uit het maken van overzichtelijke samenvattingen van dit materiaal in de vorm van tabellen en grafieken en in het globaal karakteriseren van de waarnemingen in de vorm van kengetallen (zoals gemiddelden en standaardafwijkingen), iets dat vooral bij omvangrijk waarnemingsmateriaal reeds van grote betekenis kan zijn. Bij de bevolkingsstatistiek b.v. vormt dit *beschrijvende* statistische werk het grootste deel van de statistische werkzaamheid. De techniek van het verzamelen van gegevens door middel van enquêtes - ter wille van een enorme tijdsbesparing - neemt daarbij de plaats in van de techniek van het opzetten van experimenten op gebieden als landbouw, biologie en techniek. Voor beide technieken zijn aparte onderdelen van de statistische theorie ontwikkeld.

LITERATUUR

ELEMENTAIRE INLEIDINGEN

- HEMELRIJK, J., *Statistiek, te pas en onpas*, Kluwer, 1972.
- MOSTELLER, F., W.H. KRUSKALL, e.a., *Statistics by Example* (4 delen), Addison-Wesley, 1973.
- WALLIS, W.A. & H.V. ROBERTS, *Statistics: A New Approach*, Free Press, 1956.
- NOETHER, G., *Introduction to Statistics: A Nonparametric Approach*, Houghton-Mifflin, 1976.
- HODGES, J.L. & E.L. LEHMANN, *Basic Concepts of Probability and Statistics*, Holden-Day, 1970.
- JONGE, de H., *Inleiding tot de medische statistiek* (2 delen), Nederlands Instituut voor Praeventieve Geneeskunde, 1963 (deel 1), 1964 (deel 2).

DIXON, W.J. & F.J. MASSEY, *Introduction to Statistical Analysis*, McGraw-Hill, 1969.

SNEDECOR, G.W. & W.G. COCHRAN, *Statistical Methods*, Iowa State University Press, 1967.

WISKUNDIGE LEERBOEKEN VOOR STUDENTEN

SVERDRUP, E., *Laws and Chance Variations* (2 delen), North-Holland, 1967.

HOGG, R.V. & A.T. CRAIG, *Introduction to Mathematical Statistics*, MacMillan, 1970.

MOOD, A.M., F.A. GRAYBILL & D. BOES, *Introduction to the Theory of Statistics*, McGraw-Hill, 1974.

LINDGREN, B.W., *Statistical Theory*, MacMillan, 1976.

WISKUNDIGE LEERBOEKEN VOOR GEVORDERDEN

LEHMANN, E.L., *Testing Statistical Hypotheses*, Wiley, 1959.

FERGUSON, T.S., *Mathematical Statistics: A Decision Theoretic Approach*, Academic Press, 1967.

RAO, C.R., *Linear Statistical Inference and its Applications*, Wiley, 1973.

ZACKS, S., *The Theory of Statistical Inference*, Wiley, 1971.

SCHMETTERER, L., *Introduction to Mathematical Statistics*, Springer-Verlag, 1974.

WITTING, H., *Mathematische Statistik*, Teubner, 1966.

WITTING, H. & G. NÖLLE, *Angewandte Mathematische Statistik*, Teubner, 1970.

BARRA, J.R., *Notions Fondamentales de Statistique Mathématique*, Dunod, 1971.

NIET-PARAMETRISCHE STATISTIEK

BRADLEY, J.V., *Distribution-Free Statistical Tests*, Prentice-Hall, 1968.

KENDALL, M.G., *Rank Correlation Methods*, Griffin, 1962.

NOETHER, G., *Elements of Nonparametric Statistics*, Wiley, 1967.

- GIBBONS, J.D., *Nonparametric Methods for Quantitative Analysis*, Holt, Rinehart and Winston, 1976.
- CONNOVER, W.J., *Practical Nonparametric Statistics*, Wiley, 1971.
- HOLLANDER, M. & D.A. WOLFE, *Nonparametric Statistical Methods*, Wiley, 1973.
- HÁJEK, J., *Nonparametric Statistics*, Holden-Day, 1969.
- GIBBONS, J.D., *Nonparametric Statistical Analysis*, McGraw-Hill, 1971.
- LEHMANN, E.L., *Nonparametrics: Statistical Methods Based on Ranks*, Holden-Day, 1975.
- HÁJEK, J. & Z. ŠIDÁK, *Theory of Rank Tests*, Academic Press, 1967.
- PURI, M.L. & P.K. SEN, *Nonparametric Methods in Multivariate Analysis*, Wiley, 1971.
- STATISTISCHE PROEFOPZETTEN, REGRESSIE-en VARIANTIE-ANALYSE
- COX, D.R., *Planning of Experiments*, Wiley, 1958.
- KEMPTHORNE, O., *The Design and Analysis of Experiments*, Wiley, 1952.
- WINER, B.J., *Statistical Principles of Experimental Design*, McGraw-Hill, 1971.
- MENDENHALL, W., *Introduction to Linear Models and the Design and Analysis of Experiments*, Wadsworth, 1968.
- JOHN, P.W.M., *Statistical Design and Analysis of Experiments*, MacMillan, 1971.
- GRAYBILL, F.A., *An Introduction to the Theory and Application of the Linear Model*, Duxbury Press, 1976.
- ACTON, F.S., *Analysis of Straight-line Data*, Wiley, 1959.
- DRAPER, M.R. & H. SMITH, *Applied Regression Analysis*, Wiley, 1966.
- DANIEL, C. & F.S. WOOD, *Fitting Equations to Data*, Wiley, 1971.
- SCHEFFÉ, H., *The Analysis of Variance*, Wiley, 1959.
- SEARLE, S.R., *Linear Models*, Wiley, 1971.

TUCKER, H., *A Graduate Course in Probability*, Academic Press, 1967.

CHUNG, K.L., *A Course in Probability Theory*, Harcourt, Brace & World, 1976.

STATISTISCHE HANDBOEKEN

KENDALL, M.G. & A. STUART, *The Advanced Theory of Statistics* (3 delen), Griffin, 1971 (deel 1), 1973 (deel 2), 1976 (deel 3).

JOHNSON, N.L. & S. KOTZ, *Distributions in Statistics: Discrete Distributions*, (1969), *Continuous Univariate Distributions 1*, (1970), *Continuous Univariate Distributions 2*, (1970), *Continuous Multivariate Distributions*, (1972), Wiley.

VRAAGSTUKKEN VERZAMELINGEN

HEMELRIJK, J. & D. WABEKE, *Elementaire statistische opgaven met uitgewerkte oplossingen*, Noorduijn en Zoon, 1957.

RAHMAN, N.A., *Exercises in Probability and Statistics*, Griffin, 1967.

BARRA, J.R. & A. BAILLE, *Problèmes de Statistique Mathématique*, Dunod, 1969.

TABELLEN VERZAMELINGEN

RAND CORPORATION, *A Million Random Digits with 100,000 Normal Deviates*, Free Press, 1955.

OWEN, D.B., *Handbook of Statistical Tables*, Addison-Wesley, 1962.

PEARSON, E.S. & H.O. HARTLEY, *Biometrika Tables for Statisticians* (2 delen), Cambridge University Press, 1970 (deel 1), 1972 (deel 2).

FISHER, R.A. & F. YATES, *Statistical Tables for Biological, Agricultural, and Medical Research*, Oliver and Boyd, 1957.

INSTITUTE OF MATHEMATICAL STATISTICS, *Selected Tables in Mathematical Statistics* (3 delen), American Mathematical Society, 1973 (deel 1), 1974 (deel 2), 1975 (deel 3).

BEYER, W.H. (Ed.), *Handbook of Tables for Probability and Statistics*,
Chemical Rubber Co., 1966.

KRES, H., *Statistische Tafeln zur Multivariaten Analysis*, Springer-Verlag,
1975.

GREENWOOD, J.A. & H.O. HARTLEY, *Guide to Tables in Mathematical Statistics*,
Princeton University Press, 1962.

HOOFDSTUK II

HET EXPERIMENTELE FUNDAMENT VAN DE STATISTIEK

Het experimentele fundament van de statistiek is tot op zekere hoogte hetzelfde als dat van het dobbelspel: hoewel men bij een worp met een dobbelsteen niet kan voorspellen wat de uitkomst zal zijn, weet men, dat in een lange reeks worpen met een zorgvuldig gemaakte dobbelsteen alle 6 mogelijke uitkomsten ongeveer even vaak voorkomen. (Men is gewend dit feit in het kort uit te drukken door te zeggen, dat alle 6 mogelijke uitkomsten dezelfde kans bezitten.) Op grond van deze (experimenteel vastgestelde) gelijkwaardigheid der 6 mogelijke uitkomsten kan men dan langs volkomen elementaire weg van allerlei dobbelspelen de eigenschappen berekenen, b.v.: zal het vaker, even vaak of minder vaak gelukken in 4 worpen met één dobbelsteen minstens éénmaal 6 te werpen dan in 24 worpen met twee dobbelstenen dubbel 6? Dit soort van vraagstukken heeft in de 17e eeuw in Frankrijk geleid tot het ontstaan van de waarschijnlijkheidsrekening (CHEVALLIER DE MÉRÉ, de speler; PASCAL en FERMAT, de wiskundigen).

Een soortgelijke regelmaat als bij geluksspelen vindt men echter ook op andere terreinen. F.W.J. SÜSSKIND (1746) en A. QUETELET (1827) hebben reeds gewezen op het werkwaardige verschijnsel, dat de frequentie per jaar per 1000 inwoners van geboorte, sterfte en huwelijk, van moord en andere misdaden en ook van andere maatschappelijke verschijnselen vrij nauwkeurig constant blijft in de tijd, zolang er althans geen veranderingen in de maatschappelijke toestanden optreden. Toch kan men doorgaans van de mensen individueel niet voorspellen of ze binnen een jaar zullen sterven, een moord zullen begaan, etc. Men noemt het beschreven verschijnsel de *experimentele wet der grote aantallen*. Een iets nauwkeuriger formulering is de volgende.

Laat E een experiment voorstellen, dat een aantal (N) malen wordt

herhaald. Laat S één (willekeurige) der mogelijke uitkomsten van E voorstellen en $n(S)$ het aantal malen, dat S bij de N uitvoeringen van E , inderdaad optreedt. Dan wordt $n(S)/N$ het frequentiequotiënt ^{*)} van S in deze reeks experimenten genoemd. Verricht men verschillende van deze reeksen van experimenten, met dezelfde E , S en N , dan zullen in de regel verschillende waarden van $f_q(S)$ gevonden worden. De experimentele wet van de grote aantallen zegt nu, dat voor een grote klasse van experimenten E geldt, dat deze verschillen tussen de voor $f_q(S)$ gevonden waarden voor grote waarden van N zeer klein worden.

De term "experiment" vatte men hierbij zeer ruim op. Bij de bovenstaande voorbeelden valt daar niet alleen een worp met een dobbelsteen onder, maar ook het verrichten van een waarneming, b.v. betrekking hebbende op het al of niet in leven zijn van een bepaalde persoon op een bepaalde datum.

De kansrekening en de statistiek houden zich nu bezig met al die verschijnselen, waarvoor de experimentele wet der grote aantallen geldt of geacht wordt te gelden. De kansrekening is de zuiver wiskundige kern, die los van problemen van toepassing en verificatie op axiomatische wijze wordt opgebouwd en de statistiek houdt zich bezig met het opbouwen van begrippen en theoriën, die de brug moeten slaan tussen de praktische problemen en deze theorie.

Een praktisch voorbeeld van de experimentele wet der grote aantallen vindt men in de onderstaande tabellen, die betrekking hebben op het geslacht van in 1954 in 6 wijken van Amsterdam geboren kinderen.

Tabel I
Aantallen jongens- en meisjesgeboorten in 6 wijken in Amsterdam (1954)

wijk nr	aantal geboren			fq (jongens)
	jongens	meisjes	totaal (N)	
1	23	35	58	0,40
2	300	269	569	0,53
3	277	272	549	0,50
4	25	27	52	0,48
5	281	289	570	0,49
6	68	61	129	0,53

*)

afkorting: f_q , meervoud f_{qn} ; notatie $f_q(S)$.

De f_{qn} der jongensgeboorten vertonen in tabel I reeds ongeveer dezelfde orde van grootte, terwijl het f_q , dat het meest van de overige afwijkt, optreedt bij een kleine waarde van N . Vergroten wij de waarden van N nu door de wijken achtereenvolgens bij elkaar te nemen, dan worden de verschillen reeds vanaf de tweede regel veel geringer. Dit is in tabel II gedaan.

Tabel II

Aantallen jongens- en meisjesgeboorten
in 6 wijken in Amsterdam, cumulatief

wijk nr	aantal geborenen			f_q (jongens)
	jongens	meisjes	totaal (N)	
1	23	35	58	0,40
1 en 2	323	304	627	0,52
1 t/m 3	600	576	1176	0,51
1 t/m 4	625	603	1228	0,51
1 t/m 5	906	892	1798	0,50
1 t/m 6	974	953	1927	0,51

Hierbij valt nog op te merken, dat de f_{qn} , die in de tabellen slechts in twee decimalen zijn opgegeven, wel "nauwkeuriger" berekend kunnen worden, doch dat het - juist vanwege de statistische variabiliteit - zinloos is dit te doen. Tegen deze regel, die betrekking heeft op alle verschijnselen, die statistische variabiliteit vertonen, wordt nogal eens gezondigd, waardoor dan een onverantwoorde suggestie van precisie ontstaat.

In aansluiting hierop doet zich de vraag voor, of wij het f_q , verkregen uit alle gegevens tezamen, eigenlijk niet op 0,50 af zouden moeten ronden, d.w.z. of de tweede decimaal in dit resultaat "nog wel zin heeft". Deze vage uitdrukking kan men vervangen door de iets suggestievere: "Wijst de uitkomst 0,51 erop, dat er in Amsterdam systematisch meer jongens dan meisjes geboren worden, of is de afwijking van 0,50 slechts een toevallige?" Vraagt men zich nu af, wat men met de termen "systematisch" en "toevallig" bedoelt, dan kan het antwoord hierop als volgt luiden. Indien niet alleen in het onderzochte jaar, maar ook in voorafgaande en volgende jaren bijna altijd meer jongens dan meisjes geboren worden, spreken wij van een systematische afwijking van 0,50. In dat geval kan men dus ook

voorspellen, dat er in een komend jaar meer jongens dan meisjes geboren zullen worden en die voorspelling zal dan in de regel uitkomen. Is dit niet zo, dus is er geen systematisch verschil tussen het aantal jongens- en meisjesgeboorten, dan noemt men de gevonden afwijking van 0,50 een toevallige afwijking. Een nadere precisering van deze, ook in bovenstaande terminologie nog verre van exacte, formulering, benevens de ontwikkeling van methoden, om op grond van gegevens van de in de tabellen I en II vermelde aard uit te maken in welk geval men verkeert, is nu juist de taak van de statistiek. Daarbij wordt dan gebruik gemaakt van een exact mathematisch model, dat een vereenvoudiging en schematisering van het onderzochte probleem inhoudt en waarin zowel voor de "toevallige" als voor de "systematische" verschillen exacte begrippen worden ingevoerd. Waargenomen verschillen, zoals hier het verschil tussen 0,50 en 0,51, worden in een dergelijk model gesplitst in twee delen, n.l. een systematisch en een toevallig (de statistische vakterm is *stochastisch*) verschil, terwijl dan op grond van de gevonden aantallen getoetst kan worden of het systematische deel wellicht gelijk aan 0 is. Tevens worden de voorspellingsmogelijkheden uitgewerkt en gepreciseerd.

HOOFDSTUK III

EIGENSCHAPPEN VAN FREQUENTIEQUOTIËNTEN

Statistiek en kansrekening houden zich, zoals in het voorafgaande betoogd is, in de eerste plaats bezig met f_{qn} en het is dan ook van belang enkele eenvoudige eigenschappen van f_{qn} af te leiden, die verderop van fundamentele betekenis zullen blijken te zijn.

Wij beschouwen daartoe een eindige verzameling Λ van N elementen λ en een aantal eigenschappen (of *kenmerken*) $A, B, C, \dots$. Ieder element bezit geen, één of meer van deze kenmerken. Uit gegeven kenmerken kunnen door negatie ($\bar{A}$ = niet A), *conjunctie* ($A \wedge B$ = A en B) en *disjunctie* ($A \vee B$ = A of/en B) en door herhaling van deze operaties nieuwe kenmerken gevormd worden, waarvoor uiteraard hetzelfde opnieuw geldt.

Voor een willekeurig kenmerk S geven wij het aantal elementen λ , dat S bezit, aan met $N(S)$. Vervolgens definiëren wij het f_q van S door

$$(3.1) \quad f_q(S) \stackrel{\text{def}}{=} \frac{N(S)}{N}$$

en het *voorwaardelijke* f_q van S , onder de voorwaarde T , waarin T een kenmerk is met $N(T) \neq 0$, door

$$(3.2) \quad f_q(S|T) \stackrel{\text{def}}{=} \frac{N(S \wedge T)}{N(T)} .$$

Nemen wij voor T het kenmerk "tot Λ te behoren", dan gaat $f_q(S|T)$ over in $f_q(S)$. Andersom kan men zeggen: de voorwaardelijke f_{qn} onder voorwaarde T worden verkregen als gewone (onvoorwaardelijke) f_{qn} indien men Λ vangt door de verzameling van die elementen λ , die het kenmerk T bezitten.

De volgende eigenschappen van f_{qn} zijn nu evident.

$$(3.3) \quad \left\{ \begin{array}{ll} \text{I} & 0 \leq fq(S) \leq 1 \quad \text{voor ieder der beschouwde kenmerken.} \\ \text{II} & fq(S) = 0 \quad \text{dan en slechts dan als geen } \lambda \in \Lambda \\ & \quad \text{het kenmerk } S \text{ bezit.} \\ \text{III} & fq(S) = 1 \quad \text{dan en slechts dan als iedere } \lambda \in \Lambda \\ & \quad \text{het kenmerk } S \text{ bezit.} \\ \text{IV} & fq(S \vee U) = fq(S) + fq(U) - fq(S \wedge U) \\ & \quad \text{voor ieder paar } S, U \text{ van beschouwde} \\ & \quad \text{kenmerken.} \end{array} \right.$$

Eigenschap IV volgt direct uit

$$N(S \wedge U) = N(S) + N(U) - N(S \vee U).$$

Deling door N geeft het gewenste resultaat.

Volgens definitie (3.2) en het daaronder vermelde gelden de eigenschappen (3.3) ook voor voorwaardelijke fqn. Dan gaat b.v. IV over in:

$$fq(S \vee U | T) = fq(S | T) + fq(U | T) - fq(S \wedge U | T).$$

Verder volgt uit (3.1) en (3.2):

$$(3.4) \quad fq(S | T) = \frac{fq(S \wedge T)}{fq(T)} \quad \text{of} \quad fq(S \wedge T) = fq(S | T)fq(T).$$

Nu kunnen op grond van (3.3) en (3.4) *alleen* ^{*)}, dus zonder terug te grijpen op de definities (3.1) en (3.2), een aantal stellingen bewezen worden, die ook weer zowel voor onvoorwaardelijke als voor voorwaardelijke fqn gelden. Bij de bewijzen van deze stellingen wordt bovendien slechts gebruik gemaakt van

$$(3.3) \quad \left\{ \begin{array}{l} \text{II} \quad fq(S) = 0 \text{ als geen } \lambda \in \Lambda \text{ het kenmerk } S \text{ bezit,} \\ \text{III} \quad fq(S) = 1 \text{ als iedere } \lambda \in \Lambda \text{ het kenmerk } S \text{ bezit,} \end{array} \right.$$

in plaats van (3.3.II) en (3.3.III).

*) En onder gebruikmaking van stellingen uit de logica met betrekking tot de conjunctie en disjunctie.

De te bewijzen stellingen gelden, zoals gezegd, ook voor voorwaardelijke fqn, d.w.z. alle stellingen blijven gelden, indien men aan alle fqn dezelfde voorwaarde toevoegt; dit kan alleen spaak lopen doordat er voorwaardelijke fqn ontstaan, die onbepaald zijn omdat geen enkele λ meer aan de bij het fq behorende voorwaarde voldoet.

STELLING 3.1 (algemene optelregel).

Zijn $s_1, s_2, \dots, s_k$ kenmerken en is

$$(3.5) \quad \bigvee_{i=1}^k s_i \stackrel{\text{def}}{=} s_1 \vee s_2 \vee \dots \vee s_k$$

en

$$(3.6) \quad \bigwedge_{i=1}^k s_i \stackrel{\text{def}}{=} s_1 \wedge s_2 \wedge \dots \wedge s_k,$$

dan is

$$(3.7) \quad \begin{aligned} \text{fq}\left(\bigvee_{i=1}^k s_i\right) &= \sum_i \text{fq}(s_i) - \sum_{i < j} \text{fq}(s_i \wedge s_j) + \sum_{i < j < h} \text{fq}(s_i \wedge s_j \wedge s_h) \\ &\quad - \dots + (-1)^{k-1} \text{fq}\left(\bigwedge_{i=1}^k s_i\right). \end{aligned}$$

BEWIJS door volledige inductie.

Voor $k = 2$ gaat (3.7) over in (3.3.IV). Voor $k + 1$ geldt, indien (3.7) voor k juist wordt ondersteld, volgens (3.3.IV):

$$\begin{aligned} \text{fq}\left(\bigvee_{i=1}^{k+1} s_i\right) &= \text{fq}\left(\bigvee_{i=1}^k s_i \vee s_{k+1}\right) = \\ &= \text{fq}\left(\bigvee_{i=1}^k s_i\right) + \text{fq}(s_{k+1}) - \text{fq}\left(\left(\bigvee_{i=1}^k s_i\right) \wedge s_{k+1}\right). \end{aligned}$$

Hierin is

$$\left(\bigvee_{i=1}^k s_i\right) \wedge s_{k+1} \equiv \bigvee_{i=1}^k (s_i \wedge s_{k+1}),$$

zodat, met behulp van (3.7) voor de waarde k , volgt

$$\begin{aligned}
fq\left(\bigvee_{i=1}^{k+1} S_i\right) &= \sum_{i=1}^{k+1} fq(S_i) - \sum_{i<j}^k fq(S_i \wedge S_j) + \dots + (-1)^{k-1} fq\left(\bigwedge_{i=1}^k S_i\right) + \\
&- \sum_{i=1}^k fq(S_i \wedge S_{k+1}) + \sum_{i<j}^k fq(S_i \wedge S_j \wedge S_{k+1}) - \dots + \\
&- (-1)^{k-1} fq\left(\bigwedge_{i=1}^k S_i \wedge S_{k+1}\right) = \sum_{i=1}^{k+1} fq(S_i) - \sum_{i<j}^{k+1} fq(S_i \wedge S_j) + \dots \\
&+ (-1)^k fq\left(\bigwedge_{i=1}^{k+1} S_i\right). \quad \square
\end{aligned}$$

STELLING 3.2 (bijzondere of exclusieve optelregel).

Vormen de kenmerken $S_1, S_2, \dots, S_k$ op Λ een exclusief systeem, d.w.z. dat géén $\lambda \in \Lambda$ meer dan één der kenmerken S_i bezit, dan geldt

$$(3.8) \quad fq\left(\bigvee_{i=1}^k S_i\right) = \sum_{i=1}^k fq(S_i).$$

BEWIJS. Daar $S_1, S_2, \dots, S_k$ een exclusief systeem vormen is er, voor iedere $i_1, i_2, \dots, i_h$ met $h \geq 2$, geen $\lambda \in \Lambda$ die het kenmerk

$$S_{i_1} \wedge S_{i_2} \wedge \dots \wedge S_{i_h}$$

bezit. Uit (3.3'.II) volgt dan dat voor iedere $i_1, i_2, \dots, i_h$ met $h \geq 2$

$$fq(S_{i_1} \wedge S_{i_2} \wedge \dots \wedge S_{i_h}) = 0.$$

Vult men dit in (3.7) in, dan vindt men (3.8). $\square$

STELLING 3.3. Vormen de kenmerken $S_1, S_2, \dots, S_k$ op Λ een categorisch systeem, d.w.z. dat iedere $\lambda \in \Lambda$ precies één der kenmerken S_i bezit, dan geldt:

$$(3.9) \quad \sum_{i=1}^k fq(S_i) = 1.$$

BEWIJS. Daar ieder categorisch systeem een exclusief systeem is, geldt volgens (3.8)

$$\sum_{i=1}^k fq(S_i) = fq\left(\bigvee_{i=1}^k S_i\right).$$

Verder bezit bij een categorisch systeem iedere $\lambda \in \Lambda$ het kenmerk $\bigvee_{i=1}^k S_i$.
 Uit (3.3'.III) volgt dan

$$\text{fq}\left(\bigvee_{i=1}^k S_i\right) = 1. \quad \square$$

BIJZONDER GEVAL.

$$(3.10) \quad \text{fq}(S) + \text{fq}(\bar{S}) = 1.$$

STELLING 3.4. Is $S_1, S_2, \dots, S_k$ een symmetrisch categorisch systeem op Λ ,
 d.w.z. een categorisch systeem met

$$(3.11) \quad \text{fq}(S_1) = \text{fq}(S_2) = \dots = \text{fq}(S_k)$$

dan is

$$(3.12) \quad \text{fq}(S_i) = \frac{1}{k} \quad (i = 1, 2, \dots, k).$$

BEWIJS. Dit volgt uit stelling 3.3. $\square$

STELLING 3.5. Zijn S en U twee willekeurige kenmerken, dan geldt

$$(3.13) \quad \text{fq}(S \wedge U) \leq \text{fq}(S)$$

en

$$(3.14) \quad \text{fq}(S \vee U) \geq \text{fq}(S).$$

BEWIJS. Uit

$$S \equiv (S \wedge U) \vee (S \wedge \bar{U})$$

en (3.8) volgt

$$\text{fq}(S) = \text{fq}(S \wedge U) + \text{fq}(S \wedge \bar{U}),$$

daar $S \wedge U$ en $S \wedge \bar{U}$ op Λ een exclusief systeem vormen. Verder is, volgens
 (3.3.1)

$$\text{fq}(S \wedge \bar{U}) \geq 0,$$

dus

$$fq(S) \geq fq(S \wedge U).$$

Verder is volgens (3.3.IV)

$$fq(S \vee U) = fq(S) + fq(U) - fq(S \wedge U).$$

Daar volgens (3.13): $fq(S \wedge U) \leq fq(U)$ is dus

$$fq(S \vee U) \geq fq(S). \quad \square$$

STELLING 3.6. (algemene vermenigvuldigingsregel)

Als

$$(3.15) \quad fq\left(\bigwedge_{i=1}^{k-1} S_i\right) > 0,$$

dan geldt

$$(3.16) \quad fq\left(\bigwedge_{i=1}^k S_i\right) = fq(S_1)fq(S_2|S_1)fq(S_3|S_1 \wedge S_2) \dots fq(S_k|\bigwedge_{i=1}^{k-1} S_i).$$

BEWIJS. Uit (3.13) volgt

$$fq(S_1) \geq fq(S_1 \wedge S_2) \geq fq(S_1 \wedge S_2 \wedge S_3) \geq \dots \geq fq\left(\bigwedge_{i=1}^{k-1} S_i\right).$$

Daar volgens (3.15) $fq\left(\bigwedge_{i=1}^{k-1} S_i\right) > 0$ bestaan dus alle voorwaardelijke fqn in (3.16).

De stelling wordt nu bewezen met volledige inductie. Voor $k = 2$ gaat de stelling over in (3.4)

$$fq(S_1 \wedge S_2) = fq(S_1)fq(S_2|S_1),$$

als $fq(S_1) > 0$. Onderstel nu de stelling juist voor de waarde $k - 1$, dan volgt de stelling voor k uit

$$fq\left(\bigwedge_{i=1}^k S_i\right) = fq\left(\bigwedge_{i=1}^{k-1} S_i \wedge S_k\right) = fq\left(\bigwedge_{i=1}^{k-1} S_i\right)fq(S_k|\bigwedge_{i=1}^{k-1} S_i),$$

als

$$fq\left(\bigwedge_{i=1}^{k-1} S_i\right) > 0,$$

daar

$$fq\left(\bigwedge_{i=1}^{k-2} S_i\right) \geq fq\left(\bigwedge_{i=1}^{k-1} S_i\right). \quad \square$$

STELLING 3.7. Is $T_1, T_2, \dots, T_k$ een kategorisch systeem op Λ met $f_q(T_i) > 0$ voor alle i en is S een willekeurig kenmerk, dan geldt

$$(3.17) \quad f_q(S) = \sum_{i=1}^k f_q(S|T_i) f_q(T_i).$$

BEWIJS. Daar $T_1, T_2, \dots, T_k$ een kategorisch systeem op Λ vormen, vormen $S \wedge T_i$ voor $i = 1, 2, \dots, k$ een exclusief systeem op Λ en

$$S \equiv \bigvee_{i=1}^k (S \wedge T_i).$$

Uit (3.8) volgt dan

$$f_q(S) = f_q\left(\bigvee_{i=1}^k (S \wedge T_i)\right) = \sum_{i=1}^k f_q(S \wedge T_i).$$

De stelling volgt dan uit

$$f_q(S \wedge T_i) = f_q(S|T_i) f_q(T_i),$$

als $f_q(T_i) > 0$. $\square$

STELLING 3.8 (Stelling van BAYES).

Is $T_1, T_2, \dots, T_k$ een kategorisch systeem op Λ en S een willekeurig kenmerk met $f_q(S) > 0$ en $f_q(T_i) > 0$ voor iedere i , dan geldt

$$(3.18) \quad f_q(T_i|S) = \frac{f_q(S|T_i) f_q(T_i)}{\sum_{i=1}^k f_q(S|T_i) f_q(T_i)} \quad (i = 1, 2, \dots, k).$$

BEWIJS. Volgens (3.4) is

$$f_q(S \wedge T_i) = f_q(S) f_q(T_i|S) = f_q(T_i) f_q(S|T_i) \quad (i = 1, 2, \dots, k),$$

dus

$$f_q(T_i|S) = \frac{f_q(S|T_i) f_q(T_i)}{f_q(S)}.$$

De stelling volgt dan uit (3.17). $\square$

OPMERKING. De stellingen 3.7 en 3.8 gelden ook als $T_1, T_2, \dots, T_k$ alleen categorisch zijn met betrekking tot S , d.w.z. op de deelverzameling van Λ bestaande uit alle $\lambda \in \Lambda$, die S bezitten. Ook dan vormen n.l. $S \wedge T_i$ voor $i = 1, 2, \dots, k$ een exclusief systeem op Λ en is

$$S \equiv \bigvee_{i=1}^k (S \wedge T_i).$$

BIJ HOOFDSTUK III: Voorbeeld ter verduidelijking van de begrippen
voorwaardelijke en gewone frequenties.

Gegevens betreffende 928 Nederlandse mannen:

kleur haar \ kleur ogen	rood	blond	bruin	zwart	totaal
bruin	2	7	167	35	211
intermediair	1	36	164	8	209
blauw	9	184	307	8	508
totaal	12	227	638	51	928

Onvoorwaardelijke frequenties: $f_q(\text{bruine ogen}) = \frac{211}{928} = 0,227$

$f_q(\text{blond haar}) = \frac{227}{928} = 0,245$

Voorwaardelijke frequenties: $f_q(\text{br.o.} | \text{br.h.}) = \frac{167}{638} = 0,262$

$f_q(\text{bl.o.} | \text{br.h.}) = \frac{307}{638} = 0,481$

$f_q(\text{br.h.} | \text{br.o.}) = \frac{167}{211} = 0,791$

$f_q(\text{bl.h.} | \text{br.o.}) = \frac{7}{211} = 0,033$

$f_q(\text{bl.h.} | \text{bl.o.}) = \frac{184}{508} = 0,362$

$f_q(\text{br.h.} | \text{bl.o.}) = \frac{307}{508} = 0,604.$

Toepassing algemene optelregel (st. 3.1).

$$f_q(\text{br.h. of br.o.}) = \frac{211}{928} + \frac{638}{928} - \frac{167}{928} = \frac{2 + 7 + 167 + 35 + 164 + 307}{928} = 0,735.$$

Algemene vermenigvuldigingsregel (st. 3.2).

$$f_q(\text{bl.o. en bl.h.}) = f_q(\text{bl.o.})f_q(\text{bl.h.} | \text{bl.o.}) = \frac{508}{928} \cdot \frac{184}{508} = \frac{184}{928} = 0,198.$$

Exclusieve optelregel (st. 3.3).

$$f_q(\text{r.h. of bl.h.}) = \frac{12}{928} + \frac{227}{928} = 0,258.$$

$$f_q(\text{r.h. of bl.h.} | \text{bl.o.}) = f_q(\text{r.h.} | \text{bl.o.}) + f_q(\text{bl.h.} | \text{bl.o.}) = \frac{9}{508} + \frac{184}{508} = 0,380.$$

Stelling 3.6. Neem voor $T_1, \dots, T_k$ de haarkleur en voor S: blauwe ogen,
dan is $f_q(\text{bl.o.}) = f_q(\text{bl.o.} | \text{r.h.})f_q(\text{r.h.}) + \dots + f_q(\text{bl.o.} | \text{zw.h.})f_q(\text{zw.h.}) = 0,547.$

Stelling 3.7. $T_1 \dots T_k$ haarkleur, S blauwe ogen

$$fq(bl.h. | bl.o.) = \frac{fq(bl.o. | bl.h.)fq(bl.h.)}{fq(bl.o.)} = \frac{\frac{184}{227} \cdot \frac{227}{928}}{\frac{508}{928}} = 0,362.$$

HOOFDSTUK IV

HET KANSBEGRIIP

Het kansbegrip is oorspronkelijk ontwikkeld naar aanleiding van problemen omtrent "zuivere" kansspelen, zoals het dobbelspel (mits gespeeld met "zuivere" dobbelstenen). Als algemene beschrijving van één enkele maal spelen met een dergelijk spel kan het volgende gelden:

er zijn n mogelijke "elementaire" uitkomsten $A_1, A_2, \dots, A_n$, waarvan er precies één moet optreden en het spel is symmetrisch, d.w.z. hoewel de betekenis der A_i voor de regels van het spel voor verschillende i sterk kan verschillen, verandert het spel niet van eigenschappen, indien in de verlies- en winstregels de uitkomsten $A_1, A_2, \dots, A_n$ aan een willekeurige permutatie onderworpen worden. Dit is dan een gevolg van het feit, dat $A_1, A_2, \dots, A_n$ in principe alle even vaak optreden. Oorspronkelijk drukte men dit uit door te zeggen, dat $A_1, A_2, \dots, A_n$ "gelijkelijk mogelijk" zijn.

De voor deze situatie door LAPLACE gegeven definitie van een kans is nu:

De kans op een gebeurtenis S is gelijk aan het quotiënt van het aantal (elementaire) mogelijkheden, waarbij S gerealiseerd wordt en het totale aantal mogelijkheden.)*

Zo is b.v. bij één worp met een ("zuivere") dobbelsteen de kans op het werpen van een even aantal ogen gelijk aan $3/6$; de kans op het trekken van

*) LAPLACE en verschillende andere schrijvers (b.v. HUYGENS) maakten een onderscheid tussen "kans" en "waarschijnlijkheid". Zij spraken van "het aantal kansen" waar wij "het aantal mogelijkheden" zeggen en gebruikten voor een samengestelde gebeurtenis S het woord "waarschijnlijkheid" in plaats van "kans". Het korte woord "kans" wint het echter langzamerhand van het lange woord "waarschijnlijkheid" en wij laten daarom dit onderscheid tussen de beide termen vallen.

een schoppenkaart uit een ("goed geschud") kaartspel is $13/52 = 1/4$; algemeen: de kans op ieder der elementaire uitkomsten $A_1, A_2, \dots, A_n$ is $1/n$.

Deze definitie heeft twee bezwaren. In de eerste plaats is hij kennelijk circulair zolang men geen definitie van "gelijkelijk mogelijk" geeft en de boven vermelde eis van symmetrie is zonder verdere uitwerking te vaag, om daartoe voldoende geacht te worden. Ook fysische eisen van symmetrie zijn niet voldoende. Fabriceert men b.v. een zo zuiver mogelijk symmetrische munt van zeer homogeen materiaal, doch legt men deze telkens met "kruis" boven op tafel, dan kan men desgewenst de munt "zuiver" noemen, maar het zo uitgevoerde "kruis- of munt"-spel zeker niet. Men zal dus ook aan de wijze van werpen eisen moeten stellen, om tot een "zuiver" spel te komen en het is niet zo eenvoudig deze eisen te formuleren.

Een tweede bezwaar is de beperktheid van de definitie van LAPLACE. In lang niet alle situaties, waarop de statistiek van toepassing is, laten zich gelijkwaardige mogelijkheden aanwijzen, op grond waarvan de kans op een bepaalde gebeurtenis uitgerekend kan worden. VON MISES demonstreert dit aan het begrip "sterftekans": wat zijn daarbij de mogelijkheden, waaraan men gelijke kansen toe kan kennen? Hetzelfde doet zich trouwens reeds voor bij worpen met een "valse" dobbelsteen of munt, b.v. een punaise.

Wij zullen verderop zien, hoe deze bezwaren opgeheven kunnen worden.

Beschouwen wij voorlopig nog even de definitie van LAPLACE, dan is het verband met de stellingen over fqn, die in het vorige hoofdstuk gegeven zijn, duidelijk. Een verzameling Λ van N elementen λ_i ($i = 1, 2, \dots, N$), waarbij aan λ_i het kenmerk A_i wordt toegevoegd, is nu een wiskundig model van één uitvoering van het spel en de fqn op Λ zijn precies de boven gedefinieerde kansen. In plaats van de abstracte elementen λ_i kunnen ook de A_i zelf als elementen van Λ beschouwd worden.

Wij kunnen nu dan ook de in het begin van hoofdstuk II gestelde vraag trachten te beantwoorden met behulp van een stelling uit hoofdstuk III. Deze vraag luidde als volgt. "Wat zal vaker gelukken: met 4 worpen van een dobbelsteen minstens éénmaal 6 werpen of met 24 worpen met twee dobbelstenen minstens éénmaal dubbel 6?" *)

*) Daar er bij één dobbelsteen 6 mogelijkheden zijn en bij 2 stenen 36, terwijl $4:6 = 24:36$, verwachtte de steller van deze vraag (CHEVALIER DE MÉRÉ), dat de beide spelen gelijke winstkansen zouden geven. In de praktijk won hij echter met het eerste, maar verloor met het tweede. Hij concludeerde hieruit, dat de wiskunde niet deugt.

Bij de behandeling van dit vraagstuk gaan wij uit van de onderstelling, dat de dobbelstenen "zuiver" zijn, zodat het spel symmetrisch is en de kansdefinitie van LAPLACE kan worden toegepast. Geven wij de gebeurtenis "minstens éénmaal 6 bij 4 worpen" aan met S_1 en "minstens éénmaal (6,6) bij 24 worpen met twee stenen" met S_2 , dan behoeven wij dus, volgens het bovenstaande, slechts de fqn van deze twee kenmerken op de bij de spelen behorende verzameling Λ_1 , resp. Λ_2 te berekenen.

De elementaire gebeurtenissen, die de elementen van Λ_1 vormen, zijn de 6^4 mogelijke viertallen uitkomsten van 4 worpen met een dobbelsteen. Wij geven ze aan met (A_1, A_2, A_3, A_4) , waarin A_i het aantal ogen bij de 2^e worp voorstelt. De uitkomst $A_i = 6$ geven wij kortweg aan met 6_i en de uitkomst $A_i \neq 6$ met $\bar{6}_i$. Dan is

$$S_1 \equiv 6_1 \vee 6_2 \vee 6_3 \vee 6_4,$$

en

$$\bar{S}_1 \equiv \bar{6}_1 \wedge \bar{6}_2 \wedge \bar{6}_3 \wedge \bar{6}_4.$$

Volgens (3.10) is

$$fq(S_1) = 1 - fq(\bar{S}_1),$$

zodat het voldoende is $fq(\bar{S}_1)$ te berekenen. Volgens (3.16) is:

$$fq(\bar{S}_1) = fq(\bar{6}_1)fq(\bar{6}_2|\bar{6}_1)fq(\bar{6}_3|\bar{6}_1\wedge\bar{6}_2)f(\bar{6}_4|\bar{6}_1\wedge\bar{6}_2\wedge\bar{6}_3).$$

Nu is gemakkelijk na te gaan, dat de fqn van het rechterlid op Λ_1 alle gelijk aan $5/6$ zijn, zodat

$$fq(\bar{S}_1) = (5/6)^4,$$

dus

$$fq(S_1) = 1 - (5/6)^4 = 0,518$$

is. Geheel analoog kan men $fq(S_2)$ berekenen op Λ_2 , een verzameling bestaande uit de 36^{24} mogelijke 24-tallen van uitkomsten van het tweede spel. De uitkomst is

$$fq(S_2) = 1 - (35/36)^{24} = 0,491.$$

Dit zijn dus tevens de kansen volgens LAPLACE en indien men ervan uitgaat, dat bij vele malen herhalen van de spelen voor ieder daarvan alle elementaire gebeurtenissen (ongeveer) even vaak zullen optreden, dan geven deze kansen aan, hoe vaak S_1 resp. S_2 ongeveer zullen optreden; S_1 treedt dan dus vaker dan in de helft der gevallen op en S_2 minder vaak.

Men kan nu verschillende wegen volgen, om deze opzet der kansrekening tot een goed gefundeerde theorie uit te breiden. Bij één daarvan, waarvan wij slechts heuristisch de gedachtengang zullen schetsen, wordt expliciet vastgehouden aan het uitgangspunt van gelijkwaardigheid van elementaire gebeurtenissen. Een tweede, die ondanks het feit, dat daarbij niet van gelijkwaardigheid uitgegaan wordt, in feite equivalent is met deze eerste, is een axiomatische opzet van de theorie en deze methode is tegenwoordig de meest gebruikelijke, mede vanwege de eenvoudige wijze, waarop langs deze weg exactheid kan worden bereikt.

Bij beide voorstellingswijzen is het uitgangspunt van alle toepassingen, die symmetrie of gelijke kansen betekenen, dat de betrokken gebeurtenissen bij veelvuldig herhaalde uitvoering van het experiment (of bij herhaalde waarneming) ongeveer even vaak optreden. Beide leiden ook tot dezelfde interpretatie van een kans: een kans op een gebeurtenis S betekent dat het fq van S in een lange reeks waarnemingen ongeveer gelijk aan p zal zijn. Zij steunen derhalve direct op de experimentele wet der grote aantallen (hoofdstuk II) en stemmen hierin overeen met de definitie van LAPLACE.

De eerste gedachtengang zullen wij slechts aan een voorbeeld toelichten en wel aan dat van het geslacht van pasgeborenen. Dit is één van de in hoofdstuk II vermelde voorbeelden van de experimentele wet der grote aantallen, waarbij het echter op het eerste gezicht niet duidelijk is, wat de gelijkwaardige mogelijkheden zijn. Immers er worden - zoals uit vele statistieken blijkt - meer jongens dan meisjes geboren, zodat de twee mogelijke uitkomsten "jongen" en "meisje" (ook) in dit opzicht niet als gelijkwaardig beschouwd kunnen worden. De theorie zal moeten toelaten, dat de kans op een jongen groter is dan die op een meisje. Het is echter wèl mogelijk Λ samen te stellen uit de geboorten als elementen (vroegere en toekomstige inbegrepen) en deze elementen als gelijkwaardig te beschouwen. Het fq van het kenmerk "jongen" resp. "meisje" op deze Λ is dan de kans op een jongensgeboorte. Deze verzameling Λ kan dan als basis voor de verdere

ontwikkeling van op dit probleem toe te passen statistische methoden gebruikt worden en in feite kan men zodoende de gehele statistiek opbouwen zonder het kansbegrip in te voeren, daar dit nu slechts een ander woord voor f_q is. De statistiek wordt dan volledig *frequentie-quotiëntenrekening*.

Bij de *axiomatische opzet* van de kansrekening volgt men een andere weg, waarbij men juist de gelijkwaardigheid van mogelijkheden - die de principiële moeilijkheid vormt bij de definitie van LAPLACE laat vallen. Overigens is de opbouw vrijwel analoog aan die van een regelrechte f_{qn} -rekening, op één belangrijk verschil na. Indien men n.l. wil werken met waargenomen f_{qn} , moet men rekening houden met hun variabiliteit, waardoor telkens in alle formuleringen het woord "ongeveer" insluit. De relatie "is ongeveer gelijk aan" (notatie: $\approx$) is echter niet transitief ($10^6 \approx 10^6 - 1 \approx \dots \approx 1$), hetgeen het rekenen zeer bemoeilijkt.

Anderzijds hebben f_{qn} in lange reeksen waarnemingen een neiging tot stabiliteit (experimentele wet der grote aantallen) en de kansrekening verdisconteert dit feit - daarbij tegelijkertijd bovengenoemde moeilijkheid oplossend - door aan de gebeurtenis, waarvan het f_q beschouwd wordt, een (in de regel onbekend, maar constant) getal toe te voegen, dat de *kans* op (of *waarschijnlijkheid* van) die gebeurtenis heet. Deze kans is een wiskundig analogon van het in een lange reeks waarnemingen optredende f_q , waarbij - in eerste instantie - van de onzekerheid van het f_q afgezien is. Later blijkt, dat in het kansmodel het f_q als zodanig, met zijn statistische fluctuaties, op natuurlijke wijze weer ingevoerd kan worden. Het kansbegrip wordt vastgelegd door axioma's, niet door een definitie.

Deze gang van zaken, die verderop uitvoeriger wordt beschreven, is vrijwel analoog aan die bij de opbouw en toepassing van de meetkunde. Ook daar wordt eerst afgezien van verschillen, die in de praktijk steeds bestaan. Immers twee voorwerpen bezitten nooit precies gelijke afmetingen, het is zelfs nauwelijks mogelijk een definitie van "precies gelijk" te geven, die niet reeds een modelkarakter bezit. Lijnstukken van gelijke lengte en congruente lichamen zijn echter belangrijke begrippen voor de opbouw en de toepassing van de meetkunde. Ook de begrippen "punt", "lijn" en "vlak" zijn wiskundige modelbegrippen, die niet aan definities, doch slechts aan axioma's onderworpen worden en die aan de werkelijkheid ontleend worden door schematisering. De toepassing van de meetkunde maakt dan tenslotte de waarneming en beheersing van verschillen, waarvan eerst geabstraheerd is, beter mogelijk dan oorspronkelijk het geval was. Een soortgelijke kringloop

valt waar te nemen bij de statistiek.

Bij het begrip kans heeft men dus in gedachte de "constante kern" van het f_q te karakteriseren. Het ligt daarom voor de hand voor de axioma's een aantal eenvoudige eigenschappen van f_{qn} te nemen. Daarvoor blijken nu de eigenschappen (3.3) uitermate geschikt te zijn.

Aan deze eigenschappen lag een verzameling Λ ten grondslag, waarop kenmerken gedefinieerd waren. Een dergelijke verzameling hebben wij nu ook nodig en deze verzameling, die de basis van het wiskundige model vormt, is voor elk praktisch probleem verschillend. Een praktisch probleem, waarop de statistiek van toepassing is, heeft steeds betrekking op een situatie, die verschillende uitkomsten kan hebben. Voorbeelden daarvan hebben wij reeds herhaaldelijk genoemd. Deze verschillende mogelijke uitkomsten zijn nu de elementen van de basisverzameling die wij met Γ aan zullen geven. Zij behoeven nu niet meer "gelijkelijk mogelijk" te zijn. Op Γ zijn nu weer kenmerken gedefinieerd, waaruit door negatie, conjunctie en disjunctie nieuwe kenmerken afgeleid kunnen worden.

Voorbeeld: bij één worp met een dobbelsteen ("zuiver" of "vals", dat doet nu niet ter zake) bestaat Γ uit 6 elementen, de 6 zijden van de dobbelsteen. Als kenmerken treden op: de aantallen ogen 1,2,...,6, even, oneven, $\bar{6}$, etc. Deze kenmerken worden ook "eventualiteiten" genoemd, een term (ingevoerd door Prof. Dr. D. VAN DANTZIG) die nauw aansluit bij hun karakter: eventueel optredende gebeurtenissen.

Is Γ nu eindig, dan kunnen wij de eigenschappen (3.3) letterlijk overnemen als axioma's voor de kansrekening. Vaak heeft men echter met oneindige Γ te maken. Een eenvoudig voorbeeld is het werpen met een munt tot voor het eerst kruis verkregen wordt. Het aantal worpen, dat bij dit experiment optreedt, kan dan alle waarden 1,2,... ad infinitum aannemen en Γ moet dus oneindig veel elementen bevatten. Dit maakt een kleine wijziging in de eigenschappen II en III wenselijk, terwijl een 5^e axioma toegevoegd wordt, om het werken met oneindig veel mogelijkheden vlot te doen verlopen.

De kans op een eventualiteit S wordt aangegeven met $P(S)$. Aan alle op Γ gedefinieerde eventualiteiten wordt nu een constant getal als kans toegevoegd en deze kansen (die in de regel van onbekende grootte zijn) vol-
doen aan de volgende axioma's.

$$(4.1) \quad \left\{ \begin{array}{ll} \text{Axioma I.} & 0 \leq P(S) \leq 1 \text{ voor iedere beschouwde eventualiteit } S. \\ \text{Axioma II.} & P(S) = 0 \quad \text{als geen element van } \Gamma \text{ het kenmerk } S \\ & \text{bezit.} \\ \text{Axioma III.} & P(S) = 1 \quad \text{als ieder element van } \Gamma \text{ het kenmerk } S \\ & \text{bezit.} \\ \text{Axioma IV.} & P(S \vee U) = P(S) + P(U) - P(S \wedge U) \text{ voor ieder paar } S, U \text{ van} \\ & \text{beschouwde eventualiteiten.} \end{array} \right.$$

OPMERKING. De axioma's II en III worden ook vaak als volgt uitgedrukt:

$P(S) = 0$ als S onmogelijk is; $P(S) = 1$ als S zeker is.

De op grond van de eigenschappen (3.3) voor onvoorwaardelijke fqn afgeleide stellingen gelden nu dus ook voor kansen, daar bij het bewijs van de stellingen in hoofdstuk III geen gebruik werd gemaakt van (3.3.II) en (3.3.III), maar slechts van (3.3'.II) en (3.3'.III), die overeenkomen met (4.1.II) en (4.1.III).

Eén van de stellingen (zie stelling (3.2) luidt: *Vormen de eventualiteiten $S_1, S_2, \dots, S_k$ op Γ een exclusief systeem, dan geldt*

$$(4.2) \quad P\left(\bigvee_{i=1}^k S_i\right) = \sum_{i=1}^k P(S_i).$$

Deze eigenschap geldt voor iedere eindige k . Als 5^e axioma wordt nu de overeenkomstige eigenschap voor oneindige exclusieve systemen genomen:

AXIOMA V: Vormen $S_1, S_2, \dots$ een oneindig exclusief systeem op Γ , dan geldt:

$$(4.3) \quad P\left(\bigvee_{i=1}^{\infty} S_i\right) = \sum_{i=1}^{\infty} P(S_i).$$

Het axiomasysteem is nu volledig. Een verzameling Γ met eventualiteiten en bijbehorende kansen die aan de axioma's voldoen, wordt een *kansruimte* of *waarschijnlijkheidsruimte* genoemd.

Wij gaan nu over tot de invoering van *voorwaardelijke kansen*. De definitie daarvan is analoog aan (3.4): de voorwaardelijke kans, $P(S|T)$, op de eventualiteit S onder voorwaarde van het optreden van de eventualiteit T , die een kans $P(T) > 0$ bezit, wordt gedefinieerd door

$$(4.4) \quad P(S|T) \stackrel{\text{def}}{=} \frac{P(S \wedge T)}{P(T)}.$$

Het is nu niet zonder meer duidelijk, dat de axioma's ook gelden voor voorwaardelijke kansen. Bij fqn volgde dit direct uit definitie (3.2), maar nu moeten wij deze eigenschappen bewijzen. Is dit gebeurd, dan gelden ook alle in hoofdstuk III afgeleide stellingen voor voorwaardelijke kansen.

STELLING 4.1. *Indien $P(T) > 0$ geldt:*

$$(4.5) \quad \left\{ \begin{array}{ll} \text{I.} & 0 \leq P(S|T) \leq 1 \text{ voor iedere eventualiteit } S. \\ \text{II.} & P(S|T) = 0 \quad \text{als geen der elementen van } \Gamma, \text{ die het kenmerk} \\ & \quad T \text{ bezitten, ook } S \text{ bezit.} \\ \text{III.} & P(S|T) = 1 \quad \text{als ieder element, dat het kenmerk } T \text{ bezit,} \\ & \quad \text{ook } S \text{ bezit.} \\ \text{IV.} & P(S \vee U|T) = P(S|T) + P(U|T) - P(S \wedge U|T) \text{ voor ieder paar kenmerken} \\ & \quad S \text{ en } U. \end{array} \right.$$

Volgens stelling 3.5 geldt

$$(4.6) \quad P(T \wedge S) \leq P(T).$$

Dus

$$P(S|T) = \frac{P(T \wedge S)}{P(T)} \leq 1.$$

Teller en noemer zijn verder beide niet negatief.

II. Als geen der elementen van Γ zowel T als S bezit, is volgens axioma II: $P(T \wedge S) = 0$.

III. Als ieder element, dat T bezit, ook S bezit, is volgens axioma II: $P(T \wedge \bar{S}) = 0$. Uit (4.6) volgt dan: $P(T) = P(T \wedge S)$.

$$\text{IV.} \quad P((S \vee U) \wedge T) = P((S \wedge T) \vee (U \wedge T)) = P(S \wedge T) + P(U \wedge T) - P(S \wedge U \wedge T).$$

$$\begin{aligned} P(S \vee U|T) &= \frac{P((S \vee U) \wedge T)}{P(T)} = \frac{P(S \wedge T)}{P(T)} + \frac{P(U \wedge T)}{P(T)} - \frac{P(S \wedge U \wedge T)}{P(T)} = \\ &= P(S|T) + P(U|T) - P(S \wedge U|T). \quad \square \end{aligned}$$

V. Als $S_1, S_2, \dots$ een oneindig exclusief systeem op Γ vormen dan is

$$P\left(\bigvee_{i=1}^{\infty} S_i \mid T\right) = \frac{P\left(\bigvee_{i=1}^{\infty} S_i \wedge T\right)}{P(T)} = \frac{P\left(\bigvee_{i=1}^{\infty} (S_i \wedge T)\right)}{P(T)} = \frac{\sum_{i=1}^{\infty} P(S_i \wedge T)}{P(T)} = \sum_{i=1}^{\infty} P(S_i \mid T).$$

Tenslotte voeren wij nog het begrip *stochastische onafhankelijkheid* in. De eventualiteit S is per definitie stochastisch onafhankelijk van de eventualiteit T , als

$$(4.7) \quad P(S \mid T) = P(S).$$

Volgens definitie (4.4) is dan

$$(4.8) \quad P(S \wedge T) = P(S) \cdot P(T),$$

zodat dan ook geldt:

$$(4.9) \quad P(T \mid S) = P(T).$$

De relatie van stochastische onafhankelijkheid is dus wederkerig. Tevens geldt, als S en T stochastisch onafhankelijk zijn:

$$(4.10) \quad \begin{cases} P(S \mid T) = P(S \mid \bar{T}) = P(S), \\ P(\bar{S} \mid T) = P(\bar{S} \mid \bar{T}) = P(\bar{S}) \end{cases}$$

en analoog met verwisseling van S en T . Het bewijs van deze eigenschappen en van de hieronder volgende laten wij aan de lezer over.

$$(4.11) \quad \text{als } S \rightarrow T \text{ (d.w.z. als } S \text{ T impliceert), dan is } P(S) \leq P(T) \text{ en } P(S \wedge T) = P(S), P(S \vee T) = P(T),$$

$$(4.12) \quad \text{is } P(S) = 1, \text{ dan is iedere } T \text{ stochastisch onafhankelijk van } S.$$

Stochastische onafhankelijkheid voor meer dan twee eventualiteiten wordt als volgt gedefinieerd. De eventualiteiten $S_1, S_2, \dots, S_k$ (k mag ∞ zijn) zijn stochastisch onafhankelijk, indien voor ieder h -tal ($h \leq k$)

$S_{i_1}, S_{i_2}, \dots, S_{i_h}$ daaruit geldt:

$$(4.13) \quad P\left(\bigwedge_{j=1}^h S_{i_j}\right) = \prod_{j=1}^h P(S_{i_j}).$$

Dit impliceert o.a. dat ieder tweetal S_a, S_b stochastisch onafhankelijk is en ook, dat

$$(4.14) \quad P\left(\bigwedge_{i=1}^k S_i\right) = \prod_{i=1}^k P(S_i) \quad (\text{bijzondere of onafhankelijke vermenigvuldigingsregel; voor onafhankelijke eventualiteiten})$$

Paarsgewijze onafhankelijkheid noch het vervuld zijn van (4.14) garandeert dat (4.13) vervuld is. Een tegen-voorbeeld voor het tweede is als volgt te construeren. Laat Γ uit 8 elementen bestaan, aangegeven door de nummers $1, \dots, 8$ en laat deze alle een kans $\frac{1}{8}$ bezitten. Beschouw nu de volgende kenmerken:

$$A: 1 \vee 2 \vee 3 \vee 4,$$

$$B: 3 \vee 4 \vee 5 \vee 6,$$

$$C: 4 \vee 6 \vee 7 \vee 8.$$

A en C zijn dan afhankelijk, want $\frac{1}{4} = P(A)P(C) \neq P(A \wedge C) = \frac{1}{8}$. Toch is

$$P(A) \cdot P(B) \cdot P(C) = \frac{1}{2} \cdot \frac{1}{2} \cdot \frac{1}{2} = \frac{1}{8}$$

en

$$P(A \wedge B \wedge C) = \frac{1}{8}.$$

Een voorbeeld van een drietal paarsgewijze onafhankelijke eventualiteiten, waarvoor echter de vermenigvuldigingsregel niet opgaat, is een Γ met 4 elementen (a, b, c, d) , die ieder een kans $\frac{1}{4}$ bezitten en

$$A: a \vee d$$

$$B: b \vee d$$

$$C: c \vee d.$$

De verificatie van de juistheid van deze bewering laten wij aan de lezer over.

Het is van belang na te gaan, wat de praktische interpretatie van het onafhankelijkheidsbegrip is, d.w.z. wanneer wij twee of meer mogelijk optredende gebeurtenissen S en T in het model zullen voorstellen door onafhankelijke eventualiteiten. Volgens (4.10) houdt onafhankelijkheid van

S en T in, dat de kans op S niet beïnvloed wordt door het al of niet optreden van T. Of anders gezegd: dat de informatie of T wel of niet optreedt ons geen informatie verschaft over het al of niet optreden van S. Dit is echter reeds min of meer een praktische interpretatie en indien het praktisch onaannemelijk geacht wordt, dat het al of niet optreden van T dat van S beïnvloedt, kan men deze twee dus door onafhankelijke eventualiteiten voorstellen. Men kan zich daarbij uiteraard vergissen, doch dat is bij modelkeuze altijd het geval. Verkeert men in twijfel, dan late men in het model een eventuele afhankelijkheid toe. Deze kan dan, zoals wij later zullen merken, statistisch onderzocht worden.

VOORBEELDEN. Op elkaar volgende worpen met een dobbelsteen, met de nodige voorzorgen uitgevoerd, worden gewoonlijk als onafhankelijk beschouwd, daar men aanneemt, dat een dobbelsteen "geen geheugen heeft". Door verschillende waarnemers uitgevoerde wegingen van een zelfde object, zonder dat de waarnemers van elkanders uitkomsten op de hoogte zijn, zijn onafhankelijk (zouden zij elkanders uitkomsten kennen, dan is beïnvloeding, dus afhankelijkheid, van psychologische aard niet uitgesloten). Anderzijds betekent afhankelijkheid van twee eventualiteiten nog niet, dat er een direct oorzakelijk verband bestaat; er kan ook een derde verschijnsel zijn, dat de beide eerste beïnvloedt. Beschouwt men b.v. voor ieder der na-oorlogse jaren het aantal verkochte auto's en het aantal schoolkinderen, dan vertonen deze beide kenmerken een duidelijke afhankelijkheid: in de jaren met veel schoolkinderen waren er ook veel auto's. Toch schaft men gewoonlijk geen auto aan omdat men zijn kinderen naar school moet brengen; ook stuurt men niet meer kinderen naar school omdat men een auto heeft. Het is de parallele stijging met de tijd, die hier het statistische verband veroorzaakt. Over een jaar of wat kan dit er trouwens weer heel anders uitzien. Men zij dus voorzichtig met de interpretatie van afhankelijkheden.

De definitie van LAPLACE komt nu als speciaal geval te voorschijn. Als model voor een symmetrisch spel (b.v. één worp met een zuivere dobbelsteen) wordt n.l. uiteraard een Γ genomen, waarvan alle elementen (de elementaire mogelijke uitkomsten van het spel) gelijke kans bezitten. Voor n elementen dus ieder kans $\frac{1}{n}$. Uit stelling 3.2 (de bijzondere optelregel) volgt dan direct, dat de kans op een willekeurige eventualiteit in zo'n geval berekend kan worden volgens de definitie van LAPLACE.

OPMERKINGEN. De hier beschreven practische interpretatie van het kansbegrip, als een model-analogon van een f_q , wordt de *frequentie-interpretatie* genoemd. Naast deze interpretatie zijn nog verschillende andere ontwikkeld, waaronder ook van subjectivistische aard: een kans wordt dan opgevat als een "graad van redelijk geloof" in het optreden van een eventualiteit. Wij houden ons in deze cursus geheel aan de frequentie-interpretatie, die wegens zijn objectief en verifiëerbaar karakter een hechtere basis voor de toepassingen geeft dan de andere interpretaties.

LITERATUUR.

De axiomatic van de kansrekening is afkomstig van A. KOLMOGOROV, *Grundbegriffe der Wahrscheinlichkeitsrechnung*, Springer, Berlijn 1933; Chelsea Publishing Co., New York 1946.

HOOFDSTUK V

VOORBEELDEN VAN HET BEREKENEN VAN KANSEN

Ter illustratie van het gebruik van de axioma's geven wij in dit hoofdstuk een aantal elementaire voorbeelden van berekeningen van kansen op een gegeven kansveld.

VOORBEELD 1. Wij beschouwen twee onafhankelijke uitvoeringen van een experiment, waarbij een "succes" (S) of een "mislukking" (M) op kan treden. De kans op succes is beide keren gelijk aan p ($0 < p < 1$). Gevraagd wordt de voorwaardelijke kansen $P(SM|SMVMS)$ en $P(MS|SMVMS)$ te berekenen.

OPLOSSING. Wij hebben: $P(S) = p$. $P(SVM) = 1$ (ax. III), dus volgens (3.10): $P(M) = 1 - p$. Daar beide uitvoeringen van het experiment onafhankelijk zijn, is (4.8) van toepassing, dus is, met $q = 1 - p$,

$$P(SS) = p^2, \quad P(SM) = P(MS) = pq, \quad P(MM) = q^2.$$

Verder is volgens (3.8)

$$P(SMVMS) = P(SM) + P(MS) = 2pq,$$

zodat definitie (4.4) geeft:

$$P(SM|SMVMS) = \frac{P(SM \wedge (SMVMS))}{P(SMVMS)} = \frac{P(SM)}{P(SMVMS)} = \frac{pq}{2pq} = \frac{1}{2}.$$

Hetzelfde geldt voor de andere voorwaardelijke kans, dus is de uitkomst:

$$(5.1) \quad P(SM|SMVMS) = P(MS|SMVMS) = \frac{1}{2}.$$

VOORBEELD 2. Een doosje bevat n briefjes, genummerd van 1 tot n . Deze

briefjes worden, nadat zij eerst grondig dooreen geschud zijn, één voor één blindelings uit de doos genomen en in volgorde van trekking neergelegd. Hoe groot is de kans, dat zij in een bepaalde van tevoren gegeven volgorde zullen komen te liggen?

OPLOSSING. Als model nemen wij in dit geval, dat bij elke trekking ieder der nog in de doos aanwezige briefjes gelijke kans bezit om bij de eerstvolgende trekking getrokken te worden. Als de briefjes even groot, even zwaar, even glad en op dezelfde wijze gevouwen (of niet gevouwen) zijn, kortom als men de voorzorg neemt ze - afgezien van hun nummer - vrijwel identiek te maken, voldoet dit model - naar experimenteel aangetoond kan worden - goed.*) Door dit model is het kansveld volledig beschreven.

De kans, dat de briefjes b.v. in de volgorde $1, 2, 3, \dots, n$ getrokken zullen worden is nu gemakkelijk te berekenen met behulp van stelling 3.6. Geven wij het nummer van de trekking aan met een index, dan is, daar in dit geval de definitie van LAPLACE toepasbaar is:

$$P(1_1) = \frac{1}{n}, \quad P(2_2 | 1_1) = \frac{1}{n-1}, \quad P(3_3 | 1_1 \wedge 2_2) = \frac{1}{n-2},$$

enz., dus volgens stelling 3.6

$$(5.2) \quad P(1_1 \wedge 2_2 \wedge \dots \wedge n_n) = \frac{1}{n} \cdot \frac{1}{n-1} \dots \frac{1}{2} \cdot 1 = \frac{1}{n!}.$$

Ditzelfde geldt echter ook voor iedere andere volgorde, m.a.w. alle $n!$ permutaties hebben gelijke kans $\frac{1}{n!}$.

OPMERKING. Trekkingen, waarbij alle nog aanwezige briefjes (of in het algemeen elementen) gelijke kans hebben om getrokken te worden, worden *aselecte trekkingen***) genoemd. Wij hebben hier met aselecte trekkingen zonder teruglegging te maken. Wordt ieder getrokken briefje voor de uitvoering der volgende trekking weer in de doos gelegd, dan zijn het aselecte trekkingen met teruglegging; maar dan is schudden voor iedere trekking nodig.

*) Daarbij valt op te merken, dat het dan niet nodig is tussen ieder tweetal trekkingen opnieuw te schudden; éénmaal grondig schudden van tevoren is genoeg.

**) Deze term is geïntroduceerd door Prof.Dr. D. van Dantzig.

VOORBEELD 3. Uit de doos met n briefjes, genummerd van 1 tot n , worden aselekt en zonder teruglegging k ($k \leq n$) briefjes getrokken. Hoe groot is de kans dat deze de nummers $1, \dots, k$ in één of andere volgorde bevatten?

OPLOSSING. De kans, dat deze k briefjes de nummers $1, \dots, k$ in de natuurlijke volgorde dragen is, zoals direct uit het vorige voorbeeld blijkt

$$P(1_1 \wedge 2_2 \wedge \dots \wedge k_k) = \frac{1}{n} \cdot \frac{1}{n-1} \dots \frac{1}{n-k+1} = \frac{(n-k)!}{n!}.$$

Voor iedere andere gegeven volgorde van deze k nummers geldt hetzelfde en daar er slechts één volgorde uit kan komen, is volgens stelling 3.2 (de exclusieve optelregel), de kans dat deze nummers er in de één of andere volgorde uit zullen komen gelijk aan de som van deze kansen. Daar er $k!$ verschillende volgorden mogelijk zijn, is de gezochte kans gelijk aan

$$(5.3) \quad \frac{k! (n-k)!}{n!} = \binom{n}{k}^{-1}.$$

OPMERKING. Deze uitkomst geldt uiteraard niet alleen voor de nummers $1, \dots, k$, doch voor iedere k -tal. De kans op ieder k -tal is dus dezelfde, zoals uit symmetrie-overwegingen reeds zonder meer te zien is. Daar er $\binom{n}{k}$ verschillende k -tallen zijn, volgt (5.3) hieruit met behulp van stelling 3.4. Een aselekt zonder teruglegging getrokken k -tal wordt, indien de volgorde van trekking buiten beschouwing gelaten wordt, een *steekproef* (zonder teruglegging) van de *omvang* k genoemd.

VOORBEELD 4. Uit de doos met b briefjes, genummerd van 1 tot n , worden er k ($k \leq n$) aselekt en zonder teruglegging getrokken en in volgorde van trekking neergelegd. Hoe groot is de kans, dat er op ieder van h ($h \leq k$) van tevoren aangegeven plaatsen een van tevoren gegeven nummer terecht komt?

OPLOSSING. Dit vraagstuk laat zich b.v. langs de volgende directe weg oplossen. Onder de k getrokken nummers moeten zich, wil de genoemde gebeurtenis, die wij met G aangeven, kunnen optreden, de h gegeven nummers bevinden. Er zijn $\binom{n}{k}$ k -tallen, die gelijke kans bezitten (voorbeeld 3) en hieronder zijn er $\binom{n-h}{k-h}$, die de h gegeven nummers bevatten. Derhalve is de kans, dat het getrokken k -tal de h nummers bevat gelijk aan

$$\frac{\binom{n-h}{k-h}}{\binom{n}{k}} .$$

Is nu aan deze voorwaarde voldaan, dan zijn alle permutaties der getrokken lootjes even waarschijnlijk (voorbeeld 2). Er zijn $k!$ permutaties, waaronder $(k-h)!$, waarbij G optreedt. De voorwaardelijke kans op G is dus

$$\frac{(k-h)!}{k!} ,$$

zodat

$$P(G) = \frac{(k-h)!}{k!} \frac{\binom{n-h}{k-h}}{\binom{n}{k}} = \frac{(n-h)!}{n!}$$

Deze eenvoudige uitkomst doet vermoeden, dat er een eenvoudiger oplossing bestaat. Dit is inderdaad het geval. Na het trekken van de eerste k briefjes kan men n.l. rustig doorgaan en alle briefjes uit de doos trekken. Het al of niet optreden van G wordt daardoor immers niet meer beïnvloed. Nu zijn er in totaal $n!$ even waarschijnlijke permutaties en daarvan zijn er $(n-h)!$ waarbij G optreedt; waaruit de gevonden uitkomst direct volgt.

VOORBEELD 5. Uit een doos met M witte en N zwarte knikkers worden aselekt k knikkers genomen. Hoe groot is de kans, dat de k -de knikker wit is?

- A. als de trekkingen met teruglegging geschieden;
- B. als de trekkingen zonder teruglegging geschieden.

OPLOSSING. Bij trekkingen met teruglegging is de k -de trekking een aselechte trekking uit dezelfde verzameling knikkers als de eerste, zodat volgens de definitie van LAPLACE de gevraagde kans gelijk is aan

$$(5.4) \quad \frac{M}{M+N} .$$

Bij trekkingen zonder teruglegging is de samenstelling van de knikkerverzameling bij de k -de trekking niet dezelfde als bij de eerste, maar hangt van de resultaten der eerste $k-1$ trekkingen af. Toch is de uitkomst dezelfde als eerst, zoals op de volgende wijze in te zien is. Als wij na de k -de trekking doorgaan tot alle $M+N$ knikkers getrokken zijn, beïnvloeden

wij het resultaat van de k -de trekking niet. Denken wij de knikkers genummerd van $1, \dots, M+N$, dan zijn alle $(M+N)!$ permutaties der trekkingsvolgorde even waarschijnlijk, zodat alle knikkers dezelfde kans bezitten om de k -de getrokken te zijn. Er zijn dus voor de k -de trekking weer $M+N$ mogelijkheden met gelijke kansen, waaruit volgt dat de kans op een witte k -de knikker weer gelijk is aan

$$\frac{M}{M+N}.$$

OPMERKING. De berekende kans is de onvoorwaardelijke kans, dat de k -de knikker wit is. De voorwaardelijke kans hierop, bij gegeven resultaat van de eerste $k-1$ trekkingen is afhankelijk van dit resultaat en dus in het algemeen niet gelijk aan $\frac{M}{M+N}$. Het verschil tussen aselechte trekkingen met en zonder teruglegging is dus, dat in het eerste geval de trekkingen stochastisch onafhankelijk zijn (anders gezegd: de kans is bij iedere trekking opnieuw, wat ook de vorige voor resultaat geven, dezelfde), terwijl in het laatste geval de trekkingen stochastisch afhankelijk zijn, ook al is de onvoorwaardelijke kans voor alle trekkingen dezelfde.

VOORBEELD 6. Uit een doos met M witte en N zwarte knikkers worden er n aselekt met teruglegging getrokken. Hoe groot is de kans, P_x , dat er x witte knikkers onder de n getrokken voorkomen?

OPLOSSING. Geven wij "wit" en "zwart" aan met W en Z , dan zijn, bij iedere trekking, de kansen daarop

$$P(W) = \frac{M}{M+N} (= p), \quad P(Z) = \frac{N}{M+N} (= q = 1-p).$$

De trekkingen zijn onafhankelijk (zie opmerking bij voorbeeld 5, zodat wij de bijzondere productregel (4.14) kunnen toepassen. Daaruit volgt, dat de kans op een rij uitkomsten van de vorm

$W \ W \ Z \ W \ Z \ Z \ \dots \ Z$

gelijk is aan

$$p \ p \ q \ p \ q \ q \ \dots \ q,$$

dus, als er in de rij x witte en $n-x$ zwarte voorkomen, gelijk aan $p^x q^{n-x}$. Het aantal verschillende rijen met x witte en $n-x$ zwarte is gelijk aan het

aantal x -tallen plaatsen onder n plaatsen, dus $\binom{n}{x}$. Volgens de bijzondere optelregel is dus

$$(5.5) \quad P_x = \binom{n}{x} p^x q^{n-x} \quad (\text{met } \binom{n}{0} = 1).$$

OPMERKING. Dezelfde uitkomst wordt verkregen bij iedere rij van n onafhankelijke experimenten met kans p op "succes" voor ieder daarvan; x is dan het aantal successen. Zijn b.v. de knikkers in de doos niet alle even groot of even zwaar en wordt volgens een bepaalde procedure getrokken, dan is p , de kans op een witte knikker, niet noodzakelijk meer gelijk aan $\frac{M}{M+N}$, maar (5.5) blijft geldig, alleen met een andere (wellicht onbekende) waarde van p . De grootte x kan de waarden $0, 1, \dots, n$ aannemen. Daar deze $n+1$ mogelijkheden een categorisch systeem vormen, geldt volgens stelling 3.3

$$(5.6) \quad \sum_{x=0}^n P_x = \sum_{x=0}^n \binom{n}{x} p^x q^{n-x} = 1,$$

hetgeen overeenkomt met de binomiale ontwikkeling van

$$(p+q)^n.$$

VOORBEELD 7. Uit een doos met M witte en N zwarte knikkers worden er n aselekt zonder teruglegging getrokken. Hoe groot is de kans, P'_x , dat er x witte knikkers onder de n getrokkenen voorkomen?

OPLOSSING. Volgens voorbeeld 3 hebben alle $\binom{M+N}{n}$ n -tallen gelijke kans om getrokken te worden. Het aantal verschillende n -tallen, die samengesteld zijn uit x witte en $n-x$ zwarte knikkers, bedraagt

$$\binom{M}{x} \binom{N}{n-x},$$

zodat volgens de definitie van LAPLACE geldt:

$$(5.7) \quad P'_x = \frac{\binom{M}{x} \binom{N}{n-x}}{\binom{M+N}{n}}.$$

OPMERKING. De grootte x neemt nu gehele waarden aan tussen $\max(n-N, 0)$ en $\min(M, n)$, deze grenzen inbegrepen. Buiten deze grenzen is $P'_x = 0$ en het

rechterlid van (5.7) ook volgens de afspraak dat $\binom{a}{b} = 0$ als $b < 0$ of $b > a$. Derhalve geldt volgens stelling 3.3

$$(5.8) \quad \sum_x \binom{M}{x} \binom{N}{n-x} = \binom{M+N}{n}.$$

Als n klein is t.o.v. N en M is

$$(5.9) \quad p'_x \approx p_x \quad \text{met } p = \frac{M}{M+N}.$$

Dit is als volgt in te zien.

$$p'_x = \frac{M!}{x! (M-x)!} \frac{N!}{(n-x)! (N-n+x)!} \frac{n! (M+N-n)!}{(M+N)!}.$$

Hierin is, als $x \ll M^*$,

$$\frac{M!}{(M-x)!} = M(M-1) \dots (M-x+1) \approx M^x$$

en, als $n-x \ll N$,

$$\frac{N!}{(N-n+x)!} = N(N-1) \dots (N-n+x+1) \approx N^{n-x}.$$

Voor $n \ll M+N$ is verder

$$\frac{(M+N)!}{(M+N-n)!} = (M+N)(M+N-1) \dots (M+N-n+1) \approx (M+N)^n$$

en indien deze drie benaderingen ingevuld worden, komt er

$$(5.10) \quad p'_x \approx \binom{n}{x} \left(\frac{M}{N+M} \right)^x \left(\frac{N}{N+M} \right)^{n-x}.$$

Is $n \ll N, M$, dan is ook $x \ll M$, $n-x \ll N$ en $n \ll N+M$. Hieruit blijkt - wat ook zonder bewijs wel duidelijk is - dat het niet terugleggen der getrokken elementen wel verwaarloosd kan worden zolang het er slechts weinig zijn in vergelijking met de aanwezige elementen, die dezelfde kenmerken dragen.

*) " $\ll$ " staat voor veel kleiner dan.

VOORBEELD 8. Een reeks van onafhankelijke experimenten, ieder met kans p op succes, wordt voortgezet tot precies k successen verkregen zijn. Hoe groot is de kans, Q_n , dat hiervoor n uitvoeringen van het experiment nodig zijn?

OPLOSSING. De gezochte kans is gelijk aan de kans, dat bij de eerste $n-1$ experimenten $k-1$ successen verkregen worden en bij het n -de ook een succes. De kans op het laatste is p en die op het eerste wordt door (5.5) gegeven met $n-1$ in plaats van n en $k-1$ in plaats van x .

Wegens de onafhankelijkheid der experimenten moeten deze beide kansen met elkaar vermenigvuldigd worden, hetgeen leidt tot

$$(5.11) \quad Q_n = \binom{n-1}{k-1} p^k q^{n-k} \quad (q = 1 - p).$$

Verder geldt voor $k \geq 1$

$$(5.12) \quad \sum_{n=k}^{\infty} Q_n = \sum_{n=k}^{\infty} \binom{n-1}{k-1} p^k q^{n-k} = \begin{cases} 1 & \text{als } p \neq 0 \\ 0 & \text{als } p = 0. \end{cases}$$

BEWIJS. Iedere Q_n bevat een factor p^k met $k \geq 1$, dus

$$\sum_{n=k}^{\infty} Q_n = 0 \quad \text{als } p = 0.$$

Stel nu

$$S_k = \sum_{n=k}^{\infty} Q_n = \sum_{n=k}^{\infty} \binom{n-1}{k-1} p^k q^{n-k},$$

dan is

$$qS_k = \sum_{n=k}^{\infty} \binom{n-1}{k-1} p^k q^{n-k+1} = \sum_{n=k+1}^{\infty} \binom{n-2}{k-1} p^k q^{n-k},$$

dus

$$\begin{aligned} pS_k &= p^k + \sum_{n=k+1}^{\infty} \left[\binom{n-1}{k-1} - \binom{n-2}{k-1} \right] p^k q^{n-k} = p^k + \sum_{n=k+1}^{\infty} \binom{n-2}{k-2} p^k q^{n-k} = \\ &= \sum_{n=k}^{\infty} \binom{n-2}{k-2} p^k q^{n-k} = pS_{k-1}. \end{aligned}$$

Voor $p \neq 0$ geldt dus $S_k = S_{k-1}$ ($k = 1, 2, \dots$) en dus is

$$S_k = S_1 = \sum_{n=1}^{\infty} pq^{n-1} = p \cdot \frac{1}{1-q} = 1. \quad \square$$

VOORBEELD 9. *) Een dictator, die van oordeel is, dat de verhouding van het aantal jongensgeboorten tot het aantal meisjesgeboorten niet groot genoeg is, vaardigt een wet uit, die inhoudt dat in een gezin, waarin een meisje geboren is, verder geen kinderen ter wereld mogen komen. Wat is de invloed van deze maatregel, indien aangenomen mag worden dat het geslacht van een nieuwe baby onafhankelijk is van dat van eventuele vroegere baby's van het gezin, terwijl de kans op een jongensgeboorte voor alle gezinnen dezelfde is?

OPLOSSING. Daar het door deze maatregel aan gezinnen, waarin een meisje geboren is, verboden wordt nog meer kinderen te krijgen, zal het totale aantal geboorten dalen, indien daar geen compensatie tegenover wordt gesteld. Onder de gestelde omstandigheden wordt echter de verhouding tussen de aantallen jongens- en meisjesgeboorten niet beïnvloed, daar bij iedere nieuwe geboorte de kans op een jongen gelijk aan p is, in welk gezin deze geboorte ook plaatsvindt. Men kan het ook zo stellen; de ooeivaar, die de kinderen brengt, neemt telkens een baby aselekt uit een voorraad, waarin een vaste verhouding der geslachten aanwezig is. Door de invoering van de nieuwe wet brengt hij sommige baby's op andere adressen dan anders het geval zou zijn geweest, maar hij put ze op dezelfde wijze uit zijn voorraad, zodat er, afgezien van een zekere vertraging in de aflevering, niets aan de situatie verandert.

OPMERKING. Indien de kans op een jongensgeboorte van gezin tot gezin verschilt heeft de maatregel wel invloed op de geslachts-samenstelling van de bevolking. Immers dan wordt de omvang van gezinnen met een kleinere kans op jongens sterker verkleind dan die met een grotere kans op jongens, omdat in de regel in de eerstgenoemde gezinnen eerder een meisje geboren wordt dan in de laatstgenoemde.

*) Dit voorbeeld werd mij medegedeeld door Prof.Dr. N.G. DE BRUIJN. Het heeft niet direct betrekking op het berekenen van kansen, maar dient ter illustratie van het onafhankelijkheidsbegrip.

HOOFDSTUK VI

HET METEN VAN KANSSEN

In het vorige hoofdstuk zijn een aantal voorbeelden gegeven van het *berekenen* van kansen. Dit kan geschieden, zoals in voorbeeld 6, door de berekening te baseren op een gegeven kans p , die dan ook in het resultaat (5.5) voorkomt; in voorbeeld 1 is uitgegaan van een onbekende kans p , die geëlimineerd werd, zodat een bekend getal als uitkomst verkregen werd; ook bij de overige voorbeelden, behalve bij het laatste, konden de gevraagde kansen expliciet berekend worden, op grond van bepaalde gegevens, die de aard van symmetrievoorwaarden bezitten. In feite werd in al die gevallen, waarin numerieke uitkomsten verkregen zijn, uitgegaan van de gelijkheid van een aantal kansen.

Voor de praktische toepassing van de kansrekening kan men echter met de beheersing van dergelijke eenvoudige situaties niet volstaan. Algemeen geldt, dat het voor de toepasbaarheid van een axiomatisch ingevoerd maatbegrip nodig is dit ook praktisch meetbaar te maken, d.w.z. experimentele middelen aan te geven om er in praktijkproblemen een (min of meer) bepaalde waarde aan toe te kennen. Denkt men b.v. aan het begrip "lengte", dan zal men een praktisch uitvoerbare procedure aan moeten geven, die aan de volgende eisen voldoet.

- a. Men moet na kunnen gaan of twee voorwerpen even lang zijn en, zo neen, welk het langst is.
- b. Men voert een materiële standaard (standaardmeter) in.
- c. Op grond van a en b moet aan ieder voorwerp een getal als lengte toegevoegd kunnen worden.

Pas als aan deze eisen voldaan is, kan de theorie ook in toepassing gebracht worden en deze toepassing strekt zich dan ook alleen uit tot objecten, waarvoor instrumenten ontworpen zijn, waarmee aan deze eisen vol-

daan kan worden.

Hetzelfde geldt voor het kansbegrip en de theorie daarvan (de kansrekening). Het instrumentarium, dat aan de toepasbaarheid ten grondslag ligt, wordt grotendeels geleverd door de statistiek, waarbij gebruik gemaakt wordt van de kansrekening, o.a. in de vorm van eenvoudige principes, die in de behandelde voorbeelden besloten zijn.

Punt b. is daarbij het eenvoudigst, daar een absolute maat voor het kansbegrip reeds in de axioma's II en III vervat is. Wij kunnen dan ook, zoals uit voorbeeld 1 van hoofdstuk V blijkt, langs eenvoudige weg en met simpele middelen een kans van willekeurige gegeven grootte "construeren".

De punten a en c eisen meer voorbereiding. Wij zullen verderop in deze cursus de oplossing van a resp. c ontmoeten in de vorm van statistische toetsen voor de hypothesen dat twee onbekende kansen gelijk zijn resp. dat een onbekende kans een gegeven waarde bezit. De theorie der betrouwbaarheidsintervallen - ten nauwste verbonden aan de toetsingstheorie - maakt het mogelijk grenzen aan te geven voor het verschil tussen twee onbekende kansen resp. de waarde van één onbekende kans, terwijl de schattingstheorie de middelen beschrijft om een zo goed mogelijke getalwaarde ("schatting") van kansen en verschillen daarvan te verkrijgen en de nauwkeurigheid daarvan te onderzoeken. Al deze methoden berusten op de verwerking van waarnemingen, die verricht dienen te worden onder zodanige voorwaarden, dat de wiskundige modellen, waarop deze statistische methoden berusten, een getrouwe afspiegeling van de werkelijkheid vormen.

Deze methoden, die alle verderop uitvoerig ter sprake komen, vormen uiteraard slechts een klein onderdeel van de statistiek, een eerste stap op de goede weg.

HOOFDSTUK VII

STOCHASTISCHE GROOTHEDEN, KANSVERDELINGEN

In dit en het volgende hoofdstuk worden een aantal belangrijke begrippen en stellingen uit de kansrekening behandeld.

DEFINITIE 7.1. Een grootheid $\underline{x}$, die verschillende waarden aan kan nemen, wordt een *stochastische grootheid*^{*)} genoemd indien voor ieder gegeven reëel getal x de kans, dat $\underline{x}$ een waarde aanneemt, die hoogstens daaraan gelijk is, gedefinieerd is.

VOORBEELDEN. Grootheden, die in een wiskundig model door stochastische grootheden voorgesteld kunnen worden, zijn de volgende. Het aantal successen in een reeks onafhankelijke experimenten met ieder een kans p op succes. Het aantal worpen met een munt tot voor het eerst kruis verkregen wordt. Maar ook: de gewichtstoename van een proefdier tijdens een proefperiode; de dikte van de speklaag van de 5^e walvis, die op een bepaalde dag door een walvisvaarder gevangen wordt; het aantal telefoongesprekken, dat gedurende een bepaald tijdsinterval over een bepaalde lijn gevoerd wordt.

NOTATIE. Stochastische grootheden geven wij aan met onderstreepte latijnse letters: $\underline{x}, \underline{y}, \underline{z}, \underline{x}_1, \underline{x}_2, \underline{x}_3$ enz., terwijl getalwaarden, die zij kunnen aan nemen, vaak beschouwd als algebraïsche variabelen, door dezelfde letters zonder onderstreping aangegeven worden.

$$(7.1) \quad P(\underline{x} \leq x) \stackrel{\text{def}}{=} \text{de kans dat } \underline{x} \text{ een waarde } \leq \text{ aanneemt.}$$

^{*)} στοχαστικός = raden; een grootheid, naar de waarde waarvan men slechts gissen kan.

In plaats van onderstreepte letters worden vaak ook vet gedrukte letters of hoofdletters gebruikt.

BIJBEHORENDE TERMEN. De *kansverdeling* van een stochastische grootheid $\underline{x}$ is de verzameling van alle waarden x , die $\underline{x}$ aan kan nemen met de bijbehorende kansen

$$P(\underline{x} \leq x).$$

De *verdelingsfunctie* van $\underline{x}$ is de functie

$$(7.2) \quad F(x) \stackrel{\text{def}}{=} P(\underline{x} \leq x),$$

opgevat als functie van de algebraïsche variabele x . Iedere verdelingsfunctie is een functie, die monotoon niet dalend is en uitsluitend waarden in het gesloten interval $[0,1]$ aanneemt, terwijl

$$(7.3) \quad F(-\infty) = 0, \quad F(+\infty) = 1.$$

DEFINITIE 7.2. *Discrete*^{*)} kansverdelingen zijn kansverdelingen, waarbij de beschouwde stochastische grootheid een eindig of aftelbaar aantal waarden $x_1, x_2, \dots$ aan kan nemen, terwijl

$$(7.3) \quad p_i \stackrel{\text{def}}{=} P(\underline{x} = x_i) > 0, \quad \sum_i p_i = 1.$$

De bijbehorende verdelingsfuncties zijn trapfuncties:

$$(7.4) \quad F(x) = \sum_{x_i \leq x} p_i.$$

VOORBEELDEN. De *binomiale verdeling* (voorbeeld 6 van hoofdstuk V):

$$(7.5) \quad P(\underline{x} = x) = \binom{n}{x} p^x q^{n-x} \quad (x = 0, 1, \dots, n; q = 1 - p).$$

De *hypergeometrische verdeling* (voorbeeld 7 van hoofdstuk V):

$$(7.6) \quad P(\underline{x} = x) = \frac{\binom{M}{x} \binom{N}{n-x}}{\binom{M+N}{n}} \quad (x \text{ geheel tussen } \max(n-N, 0) \text{ en } \min(M, n)).$$

*) Of: *discontinue*.

De *negatief binomiale verdeling* (voorbeeld 8 van hoofdstuk V):

$$(7.7) \quad P(\underline{x} = x) = \binom{x-1}{k-1} p^k q^{x-k} \quad (x = k, k+1, \dots; q = 1 - p).$$

DEFINITIE 7.3. *Continue kansverdelingen* zijn kansverdelingen, waarbij de stochastische grootte een continuüm van waarden (bijv. alle waarden van een eindig of oneindig interval) aan kan nemen, terwijl de verdelingsfunctie $F(x)$ voor iedere x continu is en voor iedere x , op hoogstens eindig veel na, continu differentieerbaar is.

De afgeleide

$$(7.8) \quad f(x) \stackrel{\text{def}}{=} \frac{d}{dx} F(x)$$

wordt dan de *verdelingsdichtheid* genoemd. Derhalve geldt

$$(7.9) \quad F(x) = \int_{-\infty}^x f(v) dv \quad *)$$

en, zoals op grond van de axioma's gemakkelijk te bewijzen is

$$(7.10) \quad P(x_1 \leq \underline{x} \leq x_2) = F(x_2) - F(x_1) = \int_{x_1}^{x_2} f(x) dx,$$

$$(7.11) \quad \int_{-\infty}^{+\infty} f(x) dx = 1,$$

$$(7.12) \quad P(\underline{x} = x) = 0 \quad \text{voor iedere } x.$$

Deze kansverdelingen treden op als limiet van discrete kansverdelingen (zoals wij later zullen zien; zij worden daarom vaak voor benaderingsdoeleinden gebruikt) en als wiskundig model bij grootheden, die quantitatief op een bij benadering continue schaal afgelezen worden. Voorbeelden daarvan

*) Als de verdeling in werkelijkheid "begint" bij een punt $x=a$ kan men ook a als ondergrens van de integraal nemen; daar dan echter $F(x)=0$ voor $x \leq a$, dus ook $f(x)=0$ in dat gebied, is de ondergrens $-\infty$ steeds correct. Analoge overwegingen gelden voor de bovengrens, indien deze het "einde" van de verdeling overschrijdt.

zijn: de lengte van een recruit, de tijdsduur van een treinrit, de snelheid van een vliegtuig.

VOORBEELDEN. De belangrijkste continue kansverdeling is de *normale*, waarvoor de volgende formules gelden:

$$(7.13) \quad f(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad *) \quad (-\infty < x < \infty),$$

waarin $\sigma (> 0)$ en μ constanten (*parameters*) zijn. De verdelingsfunctie is

$$(7.14) \quad F(x) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{(v-\mu)^2}{2\sigma^2}} dv.$$

NOTATIE. Bezit $\underline{x}$ een normale verdeling met parameters μ en σ , dan wordt dit aangegeven met $\underline{x}$ is $N(\mu, \sigma)$ -verdeeld.

De vorm van de verdeling is in figuur 7.1 aangegeven; σ is de afstand van het punt $x = \mu$ tot de abscis van het buigpunt van $f(x)$.

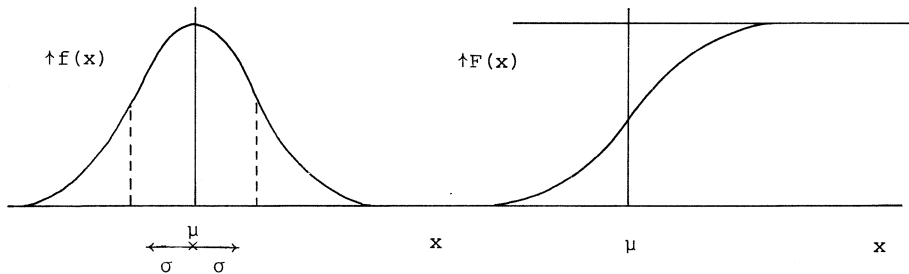


Fig. 7.1. Normale verdeling.

*) Dat voor deze functie inderdaad de relatie (7.11) geldt vindt men bewezen in §1 van de appendix.

Is $\mu = 0$ en $\sigma = 1$, dus is x $N(0,1)$ -verdeeld, dan spreekt men van de *gestandaardiseerde normale verdeling*. Voor grootheden, die deze verdeling bezitten, gebruiken wij het symbool u . De gestandaardiseerde normale verdeling is uitvoerig getabelleerd.

Een andere continue verdeling, die bijv. van belang is voor de theorie van het afronden van meet- of berekeningsresultaten, is de *homogene verdeling*, waarvoor de volgende formules gelden.

$$(7.15) \quad f(x) = \begin{cases} \frac{1}{\theta_2 - \theta_1} & \text{voor } \theta_1 \leq x \leq \theta_2, \\ 0 & \text{elders.} \end{cases}$$

$$(7.16) \quad F(x) = \begin{cases} 0 & x < \theta_1, \\ \frac{x - \theta_1}{\theta_2 - \theta_1} & \text{voor } \theta_1 \leq x \leq \theta_2, \\ 1 & \text{voor } x > \theta_2. \end{cases}$$

De "eindpunten" θ_1 en θ_2 treden hier als parameters op.*)

DEFINITIE 7.4. *Gemengde kansverdelingen* zijn kansverdelingen, die verkregen kunnen worden als gewogen gemiddelden van een discrete en een continue, d.w.z. waarvoor geldt:

$$(7.17) \quad F(x) = \lambda F_1(x) + (1-\lambda)F_2(x), \quad (0 < \lambda < 1),$$

waarin $F_1(x)$ een trapfunctie is en $F_2(x)$ continu is in x en continu differentieerbaar in alle x op hoogstens eindig veel na.

VOORBEELD. De schadeuitkering, die in één jaar op een verzekeringspolis moet geschieden. Er is een positieve kans op uitkering 0 (geen verhaalbare schade), terwijl verder een (vrijwel) continue schaal van mogelijke uitkeringen bestaat.

*) Parameters van kansverdelingen worden meestal door griekse letters aangegeven. In bepaalde gevallen, zoals bij de parameter p van een binomiale verdeling, zijn andere symbolen zozeer ingeburgerd, dat van deze regel afgeweken wordt.

OPMERKING. De mogelijkheden zijn hiermede nog niet uitgeput. Men kan continue verdelingsfuncties beschouwen, die een niet continue afgeleide bezitten of die in oneindig veel punten niet differentieerbaar zijn. Deze pathologische gevallen zijn voor de practijk van weinig of geen belang en blijven hier buiten beschouwing.

DEFINITIE 7.5. Een aantal stochastische grootheden $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ bezit een *simultane kansverdeling* als voor ieder stel reële getallen $x_1, x_2, \dots, x_n$ de (*simultane*) *verdelingsfunctie*

$$(7.18) \quad F(x_1, x_2, \dots, x_n) \stackrel{\text{def}}{=} P(\underline{x}_1 \leq x_1 \wedge \underline{x}_2 \leq x_2 \wedge \dots \wedge \underline{x}_n \leq x_n)$$

gedefinieerd is.

VOORBEELDEN. Het aantal worpen met een munt tot voor de tweede maal kruis verkregen wordt en het nummer van de worp, waarbij in die rij worpen de eerste keer kruis viel. De lichaamslengte en de borstomvang van een aselekt uit de Nederlandse mannelijke bevolking genomen man. In deze beide gevallen is $n=2$ (tweedimensionale verdelingen). Een n -dimensionale verdeling wordt bijv. verkregen indien men de temperaturen op verschillende punten van Nederland beschouwt, gemeten op een aselekt gekozen dag om 12 uur.

DEFINITIE 7.6. Een discrete n -dimensionale kansverdeling is een verdeling, waarbij in de n -dimensionale ruimte R_n met $x_1, x_2, \dots, x_n$ als coördinaten een eindig of aftelbaar aantal punten $P_j \equiv (x_{1j}, x_{2j}, \dots, x_{nj})$ ($j = 1, 2, \dots$) aanwezig zijn, waarvoor geldt

$$(7.19) \quad P(\underline{P} = P_j) = p_j > 0, \quad \sum_j p_j = 1,$$

waarin $\underline{P} \equiv (\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n)$ een *stochastisch punt* in R_n voorstelt. De bijbehorende verdeling is dan een terrasfunctie:

$$(7.20) \quad F(x_1, x_2, \dots, x_n) = \sum_{\substack{x_{ij} \leq x_i \\ (i = 1, 2, \dots, n)}} p_j$$

VOORBEELDEN. Stellen $\underline{x}_1$ resp. $\underline{x}_2$ de nummers voor van de worpen, in een rij onafhankelijke worpen met een munt, waarbij voor het eerst resp. voor de tweede maal kruis verkregen wordt, dan is, zoals gemakkelijk te zien is,

als de kans op kruis p is:

$$(7.21) \quad P(\underline{x}_1 = x_1 \wedge \underline{x}_2 = x_2) = p^2 q^{x_2 - 2} \quad (x_2 > x_1).$$

Deze kans is dus voor iedere $x_1 < x_2$ dezelfde. Dit doet vermoeden, dat bij gegeven x_2 alle mogelijke waarden van x_1 dezelfde kans zullen bezitten. Dat dit zo is kan als volgt bewezen worden (vgl. 7.7):

$$(7.22) \quad P(\underline{x}_1 = x_1 | \underline{x}_2 = x_2) = \frac{P(\underline{x}_1 = x_1 \wedge \underline{x}_2 = x_2)}{P(\underline{x}_2 = x_2)} =$$

$$= \frac{p^2 q^{x_2 - 2}}{\binom{x_2 - 1}{1} p^2 q^{x_2 - 2}} = \frac{1}{x_2 - 1}.$$

Onder de voorwaarde $\underline{x}_2 = x_2$ is $\underline{x}_1$ dus homogeen verdeeld over de getallen $1, 2, \dots, x_2 - 1$.

Laat, als tweede voorbeeld, $\underline{x}_i$ ($i = 1, 2, \dots, n$) het aantal malen kruis voorstellen bij de i -de worp uit een rij onafhankelijke worpen met een munt. Iedere $\underline{x}_i$ kan dus de waarde 0 of 1 aannemen en de kansen hierop stellen wij q resp. p . Dan geldt dus

$$P(\underline{x}_i = 1) = p, \quad P(\underline{x}_i = 0) = q \quad (i = 1, 2, \dots, n)$$

en wegens de onafhankelijkheid

$$(7.23) \quad P(\underline{x}_1 = x_1 \wedge \underline{x}_2 = x_2 \wedge \dots \wedge \underline{x}_m = x_m) = p^{\sum x_i} q^{n - \sum x_i},$$

waarin $\underline{x}_i$ ($i = 1, 2, \dots, n$) de waarde 0 of 1 bezit. Stellen wij

$$(7.24) \quad \underline{x} = \sum_{i=1}^n \underline{x}_i,$$

dan is $\underline{x}$ het aantal malen kruis bij n worpen, waarvoor dus (7.5), de binomiale verdeling, geldt.

Deze opbouw van een binomiaal verdeelde grootheid $\underline{x}$ met parameters n en p als som van een aantal onafhankelijke grootheden $\underline{x}_i$ met dezelfde p en $n=1$ zal ons later nog te pas komen.

De lezer berekene de voorwaardelijke kans

$$P(\underline{x}_1 = x_1 \wedge \underline{x}_2 = x_2 \wedge \dots \wedge \underline{x}_n = x_n | \underline{x} = x)$$

en interpreteer de verkregen uitkomst.

DEFINITIE 7.7. Een *continue n-dimensionale kansverdeling* is een kansverdeling, waarvoor de verdelingsfunctie $F(x_1, x_2, \dots, x_n)$ in elk punt $(x_1, x_2, \dots, x_n)$ continu is en waarvoor de functie

$$(7.25) \quad f(x_1, x_2, \dots, x_n) \stackrel{\text{def}}{=} \frac{\partial^n F(x_1, x_2, \dots, x_n)}{\partial x_1 \partial x_2 \dots \partial x_n}$$

overal bestaat en continu is (discontinuïteit van f in eindig veel punten of op één of meer krommen is toegestaan; een exacte formulering hiervan zou te ver voeren). De functie f heet de (*simultane*) *verdelingsdichtheid*. De omgekeerde relatie van (7.25) is

$$(7.26) \quad F(x_1, x_2, \dots, x_n) = \int_{v_1 \leq x_1} \dots \int f(v_1, v_2, \dots, v_n) dv_1 dv_2 \dots dv_n \\ (i = 1, 2, \dots, n).$$

Is G een gebied in de ruimte R_n met coördinaten $x_1, x_2, \dots, x_n$ en is $\underline{P} \equiv (x_1, x_2, \dots, x_n)$, dan geldt

$$(7.27) \quad P(\underline{P} \in G) = \int_G \dots \int f(x_1, x_2, \dots, x_n) dx_1 dx_2 \dots dx_n.$$

Het bewijs van deze relatie is elementair als G een n -dimensionaal rechthoekig parallellepipedum is. Is G van een andere vorm, dan leidt overdekking van G met inkrimpende parallellepieda en limietovergang tot het doel. Dit maattheoretische bewijs en de eisen, die aan G opgelegd moeten worden, behandelen wij hier niet.

VOORBEELDEN. Twee grootheden $\underline{x}$ en $\underline{y}$ bezitten een tweedimensionale normale verdeling, als hun verdelingsdichtheid de volgende vorm heeft:

$$(7.28) \quad f(x,y) = \frac{1}{2\pi\sigma_1\sigma_2\sqrt{1-\rho^2}} e^{-\frac{1}{2(1-\rho^2)} \left\{ \left(\frac{x-\mu_1}{\sigma_1} \right)^2 - 2\rho \frac{x-\mu_1}{\sigma_1} \frac{y-\mu_2}{\sigma_2} + \left(\frac{y-\mu_2}{\sigma_2} \right)^2 \right\}}.$$

Hierin komen 5 parameters voor: $\mu_1, \mu_2, \sigma_1, \sigma_2$ en ρ .

Deze verdeling kan vaak als model gebruikt worden indien aan één aselekt uit een grote verzameling gekozen objecten twee afmetingen waargenomen worden, bijv. lengte en breedte van een meeuwenei.

Een tweedimensionale homogene verdeling op een cirkelschijf wordt gegeven door

$$(7.29) \quad f(x,y) = \begin{cases} \frac{1}{\pi r^2} & \text{voor } x^2 + y^2 \leq r^2, \\ 0 & \text{voor } x^2 + y^2 > r^2. \end{cases}$$

Dit model treedt o.a. op bij het onderzoek van de verdeling van bacteriën of bacterie-koloniën in een Petri-schaal.

Ook in het meerdimensionale geval kunnen gemengde kansverdelingen optreden, die wij echter niet beschouwen.

Marginale verdelingen en voorwaardelijke verdelingsdichtheden.

Bezitten $x_1, x_2, \dots, x_n$ een n -dimensionale discrete verdeling (vgl. def. 7.6), dan geldt voor iedere i

$$(7.30) \quad F_i(x_i) \stackrel{\text{def}}{=} P(x_i \leq x_i) = \sum_{x_{ij} \leq x_i} p_j.$$

(Uit welke stelling volgt dit?)

$F_i(x_i)$ heet de *marginale verdelingsfunctie* van x_i . De marginale kansen worden gegeven door

$$(7.31) \quad P(x_i = x_i) = \sum_{x_{ij} = x_i} p_j$$

en analoog voor meerdimensionale marginale verdelingen.

VOORBEELD. De marginale verdeling van x_1 , behorende bij de door (7.21) gegeven verdeling, is volgens (7.31):

$$\begin{aligned} P(x_1 = x_1) &= \sum_{x_2 > x_1} p q^{x_2 - 2} = \sum_{x_2 - x_1 > 0} p q^{x_2 - x_1 - 2 + x_1} = \\ &= p q^{x_1 - 1} \sum_{i=1}^{\infty} p q^{i-1} = p q^{x_1 - 1}, \end{aligned}$$

daar $\sum_{i=1}^{\infty} p q^{i-1} = 1$. Deze marginale verdeling is dus dezelfde als (7.7) met $k = 1$. (Was dit ook direct in te zien? De lezer berekene de marginale verdeling van x_2 .)

Bezitten $x_1, x_2, \dots, x_n$ een n -dimensionale *continue* verdeling (vgl. def. 7.7), dan geldt volgens (7.27) voor de marginale verdelingsfunctie van x_i :

$$(7.32) \quad F_i(x_i) = \int_{v_i = -\infty}^{x_i} \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} f(v_1, v_2, \dots, v_n) dv_1 dv_2 \dots dv_n.$$

De *marginale verdelingsdichtheid* is dus

$$(7.33) \quad f_i(x_i) = \frac{d}{dx_i} F_i(x_i) = \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} f(v_1, v_2, \dots, x_i, \dots, v_n) \prod_{j \neq i} dv_j.$$

Meerdimensionale marginale kansverdelingen worden op analoge wijze uit de simultane verdeling afgeleid.

VOORBEELDEN. De marginale verdeling van x behorende bij (7.28) laat zich eenvoudig als volgt berekenen. Stellen wij

$$(7.34) \quad u = \frac{x - \mu_1}{\sigma_1} \quad \text{en} \quad v = \frac{y - \mu_2}{\sigma_2},$$

dan is $dx = \sigma_1 du$ en $dy = \sigma_2 dv$. Substitueren wij dit in (7.28) en stelt $::$ voor: "op een constante factor na gelijk aan", dan is dus

$$f(x,y) :: e^{-\frac{1}{2(1-\rho^2)}(u^2 - 2\rho uv + v^2)}.$$

Volgens (7.33) geldt dus voor de marginale verdelingsdichtheid $f_1(x)$ van $\underline{x}$:

$$\begin{aligned} f_1(x) &:: e^{-\frac{u^2}{2(1-\rho^2)}} \int_{-\infty}^{+\infty} e^{-\frac{1}{2(1-\rho^2)}\{(v^2 - 2\rho uv + \rho^2 u^2) - \rho^2 u^2\}} dv = \\ &= e^{-\frac{1}{2}u^2} \int_{-\infty}^{+\infty} e^{-\frac{1}{2(1-\rho^2)}(v-\rho u)^2} dv. \end{aligned}$$

Door substitutie van $w = v - \rho u$ in de integraal blijkt, dat deze - daar de grenzen door de substitutie niet veranderen - niet weer van u afhangt, zodat dus geldt

$$f_1(x) :: e^{-\frac{1}{2}u^2} = e^{-\frac{1}{2}\frac{(x-\mu_1)^2}{\sigma_1^2}}.$$

Dit betekent echter, dat de marginale verdeling van $\underline{x}$ een $N(\mu_1, \sigma_1^2)$ -verdeling is. De nog ontbrekende factor kan aangevuld worden met behulp van de voorwaarde (7.11), die uiteraard ook voor marginale verdelingen geldt. Uit formule (7.13) blijkt, dat het resultaat moet zijn

$$(7.35) \quad f_1(x) = \frac{1}{\sigma_1 \sqrt{2\pi}} e^{-\frac{1}{2}\frac{(x-\mu_1)^2}{\sigma_1^2}}.$$

Een analoge formule geldt voor $f_2(y)$, de marginale verdelingsdichtheid van $\underline{y}$.

Voorwaardelijke verdelingen kunnen, in het discrete geval direct volgens de definitie van voorwaardelijke kansen berekend worden. Een voorbeeld hiervan was (7.22). In het *continue* geval doet zich de complicatie voor, dat de voorwaarde, waaronder men werkt, zelf een kans 0 bezit, zodat de definitie niet zonder meer bruikbaar is. Bezitten $\underline{x}_1$ en $\underline{x}_2$ een continue 2-dimensionale verdeling, dan is dus $P(\underline{x}_2 = x_2) = 0$ voor iedere x_2 , zodat voorwaardelijke kansen van de vorm

$$P(\underline{x}_1 \leq x_1 | \underline{x}_2 = x_2)$$

vooralsnog niet gedefinieerd zijn. Deze kans wordt nu gedefinieerd als

$$(7.36) \quad F(x_1 | x_2) = P(\underline{x}_1 \leq x_1 | \underline{x}_2 = x_2) = \lim_{dx_2 \downarrow 0} P(\underline{x}_1 \leq x_1 | x_2 \leq \underline{x}_2 \leq x_2 + dx_2).$$

Volgens de definitie van een voorwaardelijke kans en (7.27) is dit gelijk aan

$$\lim_{dx_2 \downarrow 0} \left\{ \int_{-\infty}^{x_1} dv_1 \int_{x_2}^{x_2+dx_2} dv_2 f(v_1, v_2) \right\} / \left\{ \int_{-\infty}^{+\infty} dv_1 \int_{x_2}^{x_2+dx_2} dv_2 f(v_1, v_2) \right\}$$

en hierin is, na deling van teller en noemer door dx_2 :

$$\begin{aligned} & \lim_{dx_2 \downarrow 0} \int_{-\infty}^{x_1} dv_1 \int_{x_2}^{x_2+dx_2} dv_2 f(v_1, v_2) / dx_2 = \\ & = \frac{\partial}{\partial x_2} \int_{-\infty}^{x_1} \int_{-\infty}^{x_2} f(v_1, v_2) dv_1 dv_2 = \frac{\partial}{\partial x_2} F(x_1, x_2), \end{aligned}$$

terwijl voor de noemer geldt volgens (7.32):

$$\begin{aligned} & \lim_{dx_2 \downarrow 0} \int_{-\infty}^{+\infty} \int_{x_2}^{x_2+dx_2} dv_2 f(v_1, v_2) / dx_2 = \frac{d}{dx_2} \int_{-\infty}^{+\infty} dv_1 \int_{-\infty}^{x_2} dv_2 f(v_1, v_2) = \\ & = \frac{d}{dx_2} F_2(x_2) = f_2(x_2), \end{aligned}$$

als F_2 en f_2 de marginale verdelingsfunctie resp. -dichtheid van $\underline{x}_2$ voorstellen. Derhalve is

$$(7.37) \quad F(x_1 | x_2) = \frac{\frac{\partial}{\partial x_2} F(x_1, x_2)}{f_2(x_2)}$$

en

$$(7.38) \quad f(x_1 | x_2) = \frac{\partial}{\partial x_1} F(x_1 | x_2) = \frac{f(x_1, x_2)}{f_2(x_2)}.$$

De lezer controleer of $f(x_1|x_2)$ inderdaad een kansverdeling is, d.w.z. of $\int_{-\infty}^{+\infty} f(x_1|x_2) dx_1 = 1$ is.)

Voor meer dan 2 dimensies kan op analoge wijze de volgende formule afgeleid worden:

$$(7.39) \quad \begin{aligned} f(x_1, x_2, \dots, x_k | x_{k+1}, x_{k+2}, \dots, x_n) &= \\ &= f(x_1, x_2, \dots, x_n) \Bigg/ \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} f(x_1, x_2, \dots, x_n) dx_1 dx_2 \dots dx_k. \end{aligned}$$

VOORBEELD. Wij berekenen nu de voorwaardelijke verdelingsdichtheid $f(x|y)$ van de tweedimensionale normale verdeling (7.28). Deze is, volgens (7.38) gelijk aan

$$\frac{f(x, y)}{f_2(y)},$$

waarin, volgens (7.35),

$$f_2(y) :: e^{-\frac{1}{2} \frac{(y-\mu_2)^2}{\sigma_2^2}}.$$

Voeren wij weer de substitutie (7.34) uit, dan komt er

$$\begin{aligned} f(x|y) &:: e^{-\frac{1}{2(1-\rho^2)} (u^2 - 2\rho uv + v^2)} + \frac{1}{2} v^2 \\ &= e^{-\frac{1}{2(1-\rho^2)} (u^2 - 2\rho uv + \rho^2 v^2)} - \frac{1}{2(1-\rho^2)} (u - \rho v)^2 \\ &= e^{-\frac{1}{2(1-\rho^2)} \left(\frac{x-\mu_1}{\sigma_1} - \rho \frac{y-\mu_2}{\sigma_2} \right)^2} \\ &= e^{-\frac{\left[x - \left\{ \rho \frac{\sigma_1}{\sigma_2} (y-\mu_2) + \mu_1 \right\} \right]^2}{2(1-\rho^2)\sigma_1^2}} \\ &= e \end{aligned}$$

waarin nu y , wegens de voorwaarde $\underline{y} = y$, als een gegeven constante optreedt.

Derhalve volgt uit (7.13):

$$(7.40) \quad f(x|y) = \frac{1}{\sigma_1 \sqrt{2\pi(1-\rho^2)}} e^{-\frac{\left[x - \left\{ \rho \frac{\sigma_1}{\sigma_2} (y - \mu_2) + \mu_1 \right\} \right]^2}{2(1-\rho^2)\sigma_1^2}} .$$

HOOFDSTUK VIII

ONAFHANKELIJKHEID VAN STOCHASTISCHE GROOTHEDEN

In hoofdstuk IV is het begrip stochastische onafhankelijkheid gedefinieerd voor gebeurtenissen, door middel van de productformules (4.8) en (4.13). Aequivalent hiermee is, dat voorwaardelijke en onvoorwaardelijke kansen gelijk zijn (formule (4.7)). Wij sluiten hier bij die definities aan, waarbij dan de beschouwde gebeurtenissen zijn: $\underline{x}_1 \leq x_1$, $\underline{x}_2 \leq x_2$, enz.

DEFINITIE 8.1. Een aantal stochastische grootheden $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ zijn *onafhankelijk verdeeld* (stochastisch onafhankelijk), indien zij een simultane kansverdeling bezitten, waarvan de verdelingsfunctie voldoet aan:

$$(8.1) \quad F(\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n) = F_1(\underline{x}_1)F_2(\underline{x}_2)\dots F_n(\underline{x}_n),$$

waarin F de simultane verdelingsfunctie is en F_i ($i = 1, 2, \dots, n$) de marginale verdelingsfunctie van $\underline{x}_i$.

OPMERKING. Vergelijking (8.1) is identiek met:

$$(8.2) \quad P(\underline{x}_1 \leq x_1 \wedge \underline{x}_2 \leq x_2 \wedge \dots \wedge \underline{x}_n \leq x_n) = \prod_{i=1}^n P(\underline{x}_i \leq x_i).$$

Voor de eenvoud beschouwen wij verder in het bijzonder het geval van een *tweedimensionale verdeling* van stochastische grootheden $\underline{x}_1$ en $\underline{x}_2$. De uitbreiding tot meer dan 2 dimensies is gemakkelijk aan te geven en de bewijzen verlopen analoog. Definities (8.1) en (8.2) gaan voor $n=2$ over in

$$(8.1') \quad F(\underline{x}_1, \underline{x}_2) = F_1(\underline{x}_1)F_2(\underline{x}_2)$$

en

$$(8.2') \quad P(\underline{x}_1 \leq x_1 \wedge \underline{x}_2 \leq x_2) = P(\underline{x}_1 \leq x_1)P(\underline{x}_2 \leq x_2).$$

STELLING 8.1. Zijn $\underline{x}_1$ en $\underline{x}_2$ stochastisch onafhankelijk, dan geldt

$$(8.3) \quad P(\underline{x}_1 \leq x_1 | \underline{x}_2 \leq x_2) = P(\underline{x}_1 \leq x_1).$$

BEWIJS. Volgens de definitie van een voorwaardelijke kans, tezamen met

(8.1') geldt

$$P(\underline{x}_1 \leq x_1 | \underline{x}_2 \leq x_2) = \frac{F(\underline{x}_1, \underline{x}_2)}{F_2(\underline{x}_2)} = \frac{F_1(\underline{x}_1)F_2(\underline{x}_2)}{F_2(\underline{x}_2)} = F_1(\underline{x}_1) = P(\underline{x}_1 \leq x_1). \square$$

STELLING 8.2. Zijn $\underline{x}_1$ en $\underline{x}_2$ discreet en onafhankelijk verdeeld, dan geldt

$$(8.4) \quad \begin{cases} P(\underline{x}_1 \leq x_1 \wedge \underline{x}_2 = x_2) = P(\underline{x}_1 \leq x_1)P(\underline{x}_2 = x_2), \\ P(\underline{x}_1 = x_1 \wedge \underline{x}_2 \leq x_2) = P(\underline{x}_1 = x_1)P(\underline{x}_2 \leq x_2), \end{cases}$$

en

$$(8.5) \quad P(\underline{x}_1 = x_1 \wedge \underline{x}_2 = x_2) = P(\underline{x}_1 = x_1)P(\underline{x}_2 = x_2).$$

BEWIJS. Is $x_2' < x_2$ de grootste waarde $< x_2$, die door $\underline{x}_2$ aangenomen kan worden, dan is $P(\underline{x}_2 < x_2) = P(\underline{x}_2 \leq x_2')$ en daar (8.2') voor alle mogelijke waarden van x_1 en x_2 geldt is dus tevens

$$(8.6) \quad P(\underline{x}_1 \leq x_1 \wedge \underline{x}_2 < x_2) = P(\underline{x}_1 \leq x_1)P(\underline{x}_2 < x_2).$$

Is er geen dergelijke waarde x_2' , dan zijn linker- en rechterlid van (8.6) beide = 0, zodat de geldigheid behouden blijft.

Laat nu

A_1 resp. A_2 voorstellen $\underline{x}_1 \leq x_1$ resp. $\underline{x}_2 \leq x_2$,

B_1 resp. B_2 $\underline{x}_1 < x_1$ resp. $\underline{x}_2 < x_2$,

C_1 resp. C_2 $\underline{x}_1 = x_1$ resp. $\underline{x}_2 = x_2$.

Dan is $A_2 \equiv B_2 \vee C_2$, dus $A_1 \wedge A_2 \equiv (A_1 \wedge B_2) \vee (A_1 \wedge C_2)$, zodat

$$(8.7) \quad \begin{aligned} P(A_1 \wedge A_2) &= P(A_1)P(A_2) = P(A_1)P(B_2 \vee C_2) = \\ P(A_1)\{P(B_2) + P(C_2)\} &= P(A_1)P(B_2) + P(A_1)P(C_2). \end{aligned}$$

Anderzijds is volgens (8.6):

$$P(A_1 \wedge B_2) = P(A_1)P(B_2),$$

zodat

$$\begin{aligned} P(A_1 \wedge A_2) &= P((A_1 \wedge B_2) \vee (A_1 \wedge C_2)) = P(A_1 \wedge B_2) + P(A_1 \wedge C_2) = \\ &= P(A_1)P(B_2) + P(A_1 \wedge C_2). \end{aligned}$$

Uit deze relatie en (8.7) volgt direct, dat

$$(8.8) \quad P(A_1 \wedge C_2) = P(A_1)P(C_2),$$

hetgeen overeenkomt met de eerste der twee vergelijkingen (8.4). De tweede wordt uiteraard op dezelfde wijze bewezen.

Om (8.5) te bewijzen ontwikkelen wij nu $P(A_1 \wedge C_2)$ op analoge wijze met behulp van $A_1 \equiv B_1 \vee C_1$. Dit geeft:

$$\begin{aligned} P(A_1 \wedge C_2) &= P((B_1 \wedge C_2) \vee (C_1 \wedge C_2)) = P(B_1 \wedge C_2) + P(C_1 \wedge C_2) = \\ &= P(B_1)P(C_2) + P(C_1 \wedge C_2), \end{aligned}$$

daar (8.4) niet alleen geldt met $\underline{x}_1 \leq x_1$, maar ook met $\underline{x}_1 \leq x'_1$, als x'_1 de grootste waarde $< x_1$ is, die $\underline{x}_1$ aan kan nemen (vergelijk (8.6)). Anderzijds is volgens (8.8)

$$\begin{aligned} P(A_1 \wedge C_2) &= P(C_2)P(A_1) = P(C_2)\{P(B_1) + P(C_1)\} = \\ &= P(B_1)P(C_2) + P(C_1)P(C_2), \end{aligned}$$

wat, tezamen met de vorige relatie, leidt tot

$$P(C_1 \wedge C_2) = P(C_1)P(C_2),$$

dus tot (8.5). $\square$

STELLING 8.3. Zijn $\underline{x}_1$ en $\underline{x}_2$ continu en onafhankelijk verdeeld, dan geldt

$$(8.9) \quad f(x_1, x_2) = f_1(x_1)f_2(x_2),$$

waarin f de simultane en f_1 en f_2 de marginale verdelingsdichtheden voorstellen.

BEWIJS. Deze stelling volgt onmiddellijk uit definitie 7.7:

$$f(x_1, x_2) = -\frac{\partial^2 F(x_1, x_2)}{\partial x_1 \partial x_2} = \frac{\partial^2 \{F_1(x_1)F_2(x_2)\}}{\partial x_1 \partial x_2} = f_1(x_1)f_2(x_2). \quad \square$$

OPMERKINGEN. In hoofdstuk IV (blz. 33) is opgemerkt, dat de bijzondere productregel (4.14) niet voldoende is om stochastische onafhankelijkheid te garanderen. Daartoe was (4.13) nodig. In definitie 8.1 ligt echter een analoge relatie als (4.13) opgesloten. Immers vult men in (8.1) bijv. in $x_n = \infty$, dan gaat deze relatie over in dezelfde relatie, maar nu voor de eerste $n-1$ stochastische grootheden. Wij hebben hier dus aan (8.1) genoeg voor de definitie van stochastische onafhankelijkheid. Anderzijds is paarsgewijze onafhankelijkheid niet voldoende voor volledige onafhankelijkheid van meer dan twee stochastische grootheden. Dit is gemakkelijk door een tegenvoorbeeld aan te tonen. Op blz. 33 (hoofdstuk IV) hebben wij een voorbeeld aangetroffen van 3 eventualiteiten A, B en C, die paarsgewijze onafhankelijk waren, maar waarvoor

$$P(A \wedge B \wedge C) \neq P(A)P(B)P(C)$$

was. Definiëren wij nu $\underline{x}$, $\underline{y}$ resp. $\underline{z}$ als grootheden, die de waarde 1 aan nemen, als A, B resp. C optreden en de waarde 0 als A, B resp. C niet optreden, dan is $P(\underline{x}=1) = P(A)$, $P(\underline{y}=0) = P(\bar{B})$ enz., en onafhankelijkheid van de gebeurtenissen is equivalent met die van de stochastische grootheden. Deze laatste zijn dus ook paarsgewijze onafhankelijk, maar het drietal is dat niet, daar

$$P(A \wedge B \wedge C) = P(\underline{x}=1 \wedge \underline{y}=1 \wedge \underline{z}=1) \neq P(A)P(B)P(C) =$$

$$P(\underline{x}=1)P(\underline{y}=1)P(\underline{z}=1).$$

De relaties (8.5) en (8.9) (voor n dimensies opgeschreven) kunnen ook dienen om stochastische onafhankelijkheid te *definiëren*. Dan kan (8.1) daaruit gemakkelijk afgeleid worden. De hierboven gevolgde methode bezit

het voordeel, dat het continue en het discrete geval (en ook het hier niet behandelde gemengde geval) in één universele formule samengevat worden.

Uit de stellingen 8.2 en 8.3 volgt, dat voor onafhankelijke grootheden geldt:

$$(8.10) \quad P(\underline{x}_1 = x_1 | \underline{x}_2 = x_2) = P(\underline{x}_1 = x_1) \quad (\text{discontinue geval}),$$

en (volgens (7.38))

$$(8.11) \quad f(x_1 | x_2) = f_1(x_1) \quad (\text{continue geval}).$$

Of, algemeen gezegd, *bij stochastisch onafhankelijke grootheden is de voorwaardelijke verdeling van een aantal daarvan, onder voorwaarden, die alleen op de overige grootheden betrekking hebben, gelijk aan de marginale verdeling van de eerstgenoemde grootheden.*

VOORBEELDEN. Bij vele experimentele onderzoeken richt men de experimenten zo in, dat in het wiskundige model met onafhankelijke stochastische grootheden gewerkt kan worden. De wiskundige moeilijkheden worden daardoor sterk verminderd, omdat dan de eenvoudige relaties, die in dit hoofdstuk zijn afgeleid, gelden. Als voorbeeld van de vereenvoudigingen, die optreden, beschouwen wij de tweedimensionale normale verdeling (7.28). De marginale verdelingen van $\underline{x}$ en $\underline{y}$ zijn, volgens (7.35)

$$f_1(x) = \frac{1}{\sigma_1 \sqrt{2\pi}} e^{-\frac{1}{2} \frac{(x-\mu_1)^2}{\sigma_1^2}} \quad \text{en} \quad f_2(y) = \frac{1}{\sigma_2 \sqrt{2\pi}} e^{-\frac{1}{2} \frac{(y-\mu_2)^2}{\sigma_2^2}}.$$

Zijn $\underline{x}$ en $\underline{y}$ stochastisch onafhankelijk, dan geldt dus volgens (8.9)

$$(8.12) \quad f(x, y) = \frac{1}{2\pi\sigma_1\sigma_2} e^{-\frac{1}{2} \left\{ \frac{(x-\mu_1)^2}{\sigma_1^2} + \frac{(y-\mu_2)^2}{\sigma_2^2} \right\}},$$

hetgeen overeenkomt met (7.28) wanneer daarin $\rho = 0$ gesteld wordt. De simultane verdelingsdichtheid is nu dus veel eenvoudiger dan in het algemene geval: ten eerste is er één parameter minder en ten tweede kunnen allerlei berekeningen gemakkelijker uitgevoerd worden, omdat $f(x, y)$ uit twee factoren bestaat, die ieder slechts één der beide grootheden bevatten.

Voor meer dan 2 grootheden is de vereenvoudiging van dezelfde aard, terwijl het effect bij het oplossen van vele problemen nog sterker is.

Dat $\underline{x}$ en $\underline{y}$, als zij een simultane normale verdeling bezitten, dan en slechts dan onafhankelijk verdeeld zijn als $\rho = 0$, volgt ook uit de formules (7.35) en (7.40), die de marginale en de voorwaardelijke verdelingen van $\underline{x}$ voorstellen. Deze zijn dan en slechts dan identiek, als $\rho = 0$ is.

Een tweede voorbeeld van stochastisch onafhankelijke grootheden verkrijgt men door de resultaten te beschouwen van n op dezelfde wijze uitgevoerde worpen met een dobbelsteen. In dat geval kan men de worp-resultaten aangeven met $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ en het ligt voor de hand, daar een dobbelsteen geen geheugen heeft, als modelveronderstelling in te voeren, dat deze grootheden stochastisch onafhankelijk zijn en alle dezelfde (marginale) verdeling bezitten. Wij merken hierbij terloops op, dat men een dergelijke onderstelling niet klakkeloos behoeft te aanvaarden, doch experimenteel kan onderzoeken. De statistiek beschikt dan over de middelen om op grond van waarnemingsresultaten de gemaakte onderstelling op zijn juistheid te toetsen.

Indien wij voor de eenvoud veronderstellen, dat de dobbelsteen zuiver is, en indien $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ een willekeurig stel mogelijke uitkomsten is, geldt volgens (8.5)

$$(8.13) \quad P(\underline{x}_1 = \underline{x}_1 \wedge \underline{x}_2 = \underline{x}_2 \wedge \dots \wedge \underline{x}_n = \underline{x}_n) = \prod_{i=1}^n P(\underline{x}_i = \underline{x}_i) = 6^{-n}$$

en deze kans is dus voor alle mogelijke n -tallen uitkomsten dezelfde. Op grond van deze uitkomst kan dan desgewenst voor verdere berekeningen gebruik gemaakt worden van de kansdefinitie van LAPLACE, die op blz. 24 is vermeld. Hiervan is in feite in hoofdstuk IV (blz. 26) reeds impliciet gebruik gemaakt.

Als derde voorbeeld beschouwen wij een reeks worpen met een munt. Worden deze telkens op dezelfde wijze uitgevoerd, onafhankelijk van de bij vorige worpen verkregen resultaten, dan kan men de worpresultaten als stochastisch onafhankelijk beschouwen, met eenzelfde kans p op kruis voor iedere worp. Stellen nu $\underline{x}_1, \underline{x}_2$, enz. de nummers voor van de worpen, waarbij voor de eerste maal, voor de tweede maal, enz. kruis verkregen wordt, dan geldt uiteraard

$$(8.14) \quad P(\underline{x}_1 < \underline{x}_2 < \dots) = 1,$$

daar de tussen vierkante haken genoteerde eventualiteit zeker is. Deze stochastische grootheden zijn dus stochastisch afhankelijk, daar de waarden, die bijv. $\underline{x}_2$ aan kan nemen, afhangen van de waarde, die $\underline{x}_1$ aanneemt (en andersom). Beschouwt men echter de grootheden:

$$(8.15) \quad \underline{y}_1 = \underline{x}_1, \quad \underline{y}_2 = \underline{x}_2 - \underline{x}_1, \quad \underline{y}_3 = \underline{x}_3 - \underline{x}_2, \quad \text{enz.},$$

dan zijn deze grootheden weer onafhankelijk en zij bezitten alle dezelfde negatief binomiale verdeling (formule (7.7), hier met $k=1$). Dit volgt ogenblikkelijk uit de overweging, dat na iedere worp, waarbij kruis verkregen is, een "nieuwe" rij van onafhankelijke worpen "begint", waarop formule (7.7) van toepassing is. Bij verschillende berekeningen omtrent de kansverdeling van $\underline{x}_k$ (met $k > 1$) is het dan ook eenvoudiger om van de $\underline{y}$'s uit te gaan. Wij zullen daar verderop gebruik van maken.

Tenslotte een strikvraag^{*)}, die vaak niet juist wordt beantwoord. In fig. 8.1 is een 2-dimensionale discrete verdeling getekend, waarbij de kansen op de drie mogelijke uitkomsten nog onbepaald zijn gelaten.

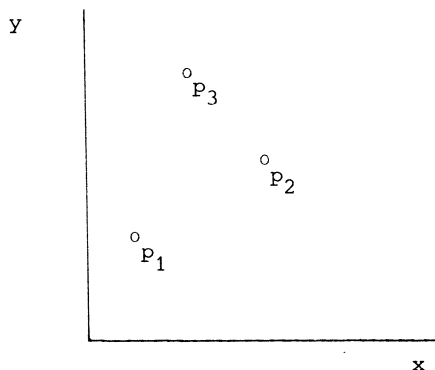


Fig.8.1. Discrete tweedimensionale verdeling

Gevraagd wordt of het mogelijk is de kansen p_1, p_2 en p_3 zo te bepalen, dat de grootheden $\underline{x}$ en $\underline{y}$ stochastisch onafhankelijk zijn.

OPGAVE.^{**)} Twee stochastische grootheden $\underline{x}$ en $\underline{y}$, die een simultane ver-

^{*)} Dit is opgave 75 van J. HEMELRIJK en DORALINE WABEKE, *Elementaire statistische opgaven met uitgewerkte oplossingen*, Noorduijn en Zoon N.V., Gorinchem 1957.

^{**)} Opgave 81 l.c.

deling bezitten, zijn onderling onafhankelijk verdeeld. Beide kunnen uitsluitend de waarden 1, 2 en 3 aannemen. De volgende kansen zijn gegeven:

x	y	$P(\underline{x} = x \wedge \underline{y} = y)$
1	3	$1/30$
2	1	$1/10$
2	2	$3/10$
3	2	$1/5$

Er zijn twee kansverdelingen, waarbij deze kansen optreden. Bereken deze.

HOOFDSTUK IX

MATHEMATISCHE VERWACHTING EN MOMENTEN

Het begrip *mathematische verwachting*, dat in dit hoofdstuk wordt besproken, is zowel voor de theorie als voor de toepassingen van groot belang. De *momenten* van een kansverdeling, die o.a. gebruikt kunnen worden om een globale beschrijving van ligging en vorm daarvan te geven, zijn speciale gevallen.

Wij behandelen eerst het *ééndimensionale geval*.

DEFINITIE 9.1. Is $\underline{x}$ een discreet verdeelde stochastische grootheid, die de waarden $x_1, x_2, \dots$ aan kan nemen, met

$$(9.1) \quad p_i = P(\underline{x}=x_i) \quad \left(\sum_i p_i = 1 \right)$$

en is $\phi(x)$ een functie van x , dan wordt de *mathematische verwachting* van $\phi(\underline{x})$ - notatie $E\phi(\underline{x})$ - gedefinieerd door

$$(9.2) \quad E\phi(\underline{x}) \stackrel{\text{def}}{=} \sum_i p_i \phi(x_i).$$

Bezit $\underline{x}$ een continue verdeling met verdelingsdichtheid $f(x)$, dan luidt de definitie

$$(9.3) \quad E\phi(\underline{x}) \stackrel{\text{def}}{=} \int_{-\infty}^{+\infty} \phi(x) f(x) dx.$$

OPMERKING. De functie $\phi(\underline{x})$ van $\underline{x}$ is, daar $\underline{x}$ stochastisch is, zelf in het algemeen ook een stochastische grootheid. De verwachting E^* is een

^{*}) Ter verkorting wordt de toevoeging "mathematische" vaak weggelaten.

"operator", die deze stochastische grootheid in een getal omzet. Immers de rechterleden van (9.2) en (9.3) zijn getallen.

Voor gemengde kansverdelingen (vergelijk blz. 51) is $E\phi(\underline{x})$ een lineaire combinatie, met coëfficiënten λ en $(1-\lambda)$ van de rechterleden van (9.2) en (9.3).

DEFINITIE 9.2. Het k -de moment ($k = 0, 1, 2, \dots$) van een stochastische grootheid $\underline{x}$ is de verwachting van $\underline{x}^k$. Notatie: μ_k of, indien nodig, $\mu_k(\underline{x})$. Dus

$$(9.4) \quad \mu_k \stackrel{\text{def}}{=} E\underline{x}^k.$$

OPMERKINGEN. Voor iedere kansverdeling geldt: $\mu_0 = 1$. Stellen wij de kansverdeling van $\underline{x}$ aanschouwelijk voor als een massaverdeling over de x -as, in het discrete geval met massapunten van massa p_i in de punten x_i en in het continue geval als een continu aangebrachte massa met dichtheid $f(x)$, dan stelt μ_0 de totale massa voor. Het eerste moment μ_1 is de abscis van het zwaartepunt van de massaverdeling en μ_2 is het traagheidsmoment om de oorsprong. De beschrijvende waarde van deze momenten voor de kansrekening is dezelfde als die in de mechanica. Aan deze analogie is het gebruik van de term "momenten" toe te schrijven.

Bij μ_1 wordt de index 1 vaak weggelaten.

VOORBEELDEN. Indien $\underline{x}$ kansen p resp. q ($p+q=1$) bezit op de waarden 1 resp. 0 (*dichotomie*), dan geldt:

$$(9.5) \quad \mu_k = p \quad \text{voor } k = 1, 2, \dots$$

Immers dan is

$$\mu_k = p \cdot 1^k + q \cdot 0^k = p.$$

Indien $\underline{x}$ een *negatief binomiale verdeling* bezit met $k=1$ (d.i. als $\underline{x}$ het aantal experimenten is tot en met het eerste succes; vergelijk voorbeeld 8 van hoofdstuk V), dan is

$$P(\underline{x}=x) = pq^{x-1} \quad (x = 1, 2, \dots)$$

en dus

$$E\underline{x} = \sum_{x=1}^{\infty} xP(\underline{x}=x) = \sum_{x=1}^{\infty} xpq^{x-1} = p \sum_{x=1}^{\infty} xq^{x-1}.$$

Nu is

$$\sum_{x=0}^{\infty} q^x = \frac{1}{1-q}$$

en nemen wij links en rechts de afgeleide, dan komt er

$$\sum_{x=1}^{\infty} xq^{x-1} = \frac{1}{(1-q)^2} = \frac{1}{p}.$$

Derhalve is

$$(9.6) \quad \mu = E\underline{x} = \frac{1}{p}.$$

Op analoge wijze valt af te leiden:

$$(9.7) \quad \mu_2 = E\underline{x}^2 = \frac{1+q}{p^2} \quad . \quad *)$$

Is $\underline{x}$ homogeen verdeeld (vergelijk pag. 51), dan is

$$E\underline{x} = \int_{\theta_1}^{\theta_2} \frac{x dx}{\theta_2 - \theta_1} = \frac{1}{2}(\theta_1 + \theta_2)$$

en

$$E\underline{x}^2 = \int_{\theta_1}^{\theta_2} \frac{x^2 dx}{\theta_2 - \theta_1} = \frac{1}{3}(\theta_1^2 + \theta_1 \theta_2 + \theta_2^2).$$

Dus

$$(9.8) \quad \mu = \frac{1}{2}(\theta_1 + \theta_2), \quad \mu_2 = \frac{1}{3}(\theta_1^2 + \theta_1 \theta_2 + \theta_2^2).$$

Is $\underline{x}$ normaal verdeeld volgens formule (7.13), dan is

$$(9.9) \quad E\underline{x} = \mu \quad \text{en} \quad E\underline{x}^2 = \sigma^2 + \mu^2,$$

waarin μ en σ de parameters uit formule (7.13) voorstellen. Voor het bewijs

*) Vergelijk voor volledige afleiding opgave 27 uit het opgavenboekje (zie voetnoot op blz. 67).

hiervan vergelijk men §1 van de appendix.

DEFINITIE 9.3. Het k -de *gereduceerde moment* ($k = 1, 2, \dots$) van een stochastische grootheid $\underline{x}$ is de verwachting van $(\underline{x} - \mu)^k$. *Notatie:* $\tilde{\mu}_k$ of, indien nodig $\tilde{\mu}_k(\underline{x})$. Verder $\tilde{x} = \underline{x} - \mu$, de *gereduceerde x*. Dus:

$$(9.10) \quad \tilde{\mu}_k = \mu_k(\tilde{x}) = E\tilde{x}^k = E(\underline{x} - \mu)^k.$$

OPMERKINGEN. Het eerste gereduceerde moment is altijd 0. Het tweede gereduceerde moment wordt de *variantie* van $\underline{x}$ genoemd en met σ^2 , zo nodig $\sigma^2(\underline{x})$, aangegeven:

$$(9.11) \quad \sigma^2(\underline{x}) = E(\underline{x} - \mu)^2 = E\tilde{x}^2 = \tilde{\mu}_2.$$

De variantie komt overeen met het traagheidsmoment om het zwaartepunt. Voor iedere kansverdeling is $\sigma^2 \geq 0$ en alleen dan 0 als $\underline{x}$ één waarde met kans 1 aanneemt. De wortel uit de variantie wordt de *spreiding* of *standaardafwijking* van $\underline{x}$ genoemd.

Het derde gereduceerde moment is 0 voor symmetrische verdelingen, maar het omgekeerde behoeft niet juist te zijn.

De berekening van de gereduceerde momenten, waarvan vooral de variantie belangrijk is, geschiedt het snelst met behulp van verderop af te leiden stellingen over de mathematische verwachting.

Het geval van *meerdere dimensionale verdelingen* behandelen wij voor het gemak in hoofdzaak voor 2 dimensies. Voor meer dan 2 verloopt het geheel analoog.

DEFINITIE 9.4. Bezitten $\underline{x}$ en $\underline{y}$ een simultane discrete verdeling, waarbij het punt $\underline{p} \equiv (\underline{x}, \underline{y})$ met kans p_j ($j = 1, 2, \dots$) in het punt $P_j \equiv (x_j, y_j)$ terecht komt en is $\phi(x, y)$ een functie van x en y , dan is de *mathematische verwachting* van $\phi(\underline{x}, \underline{y})$:

$$(9.12) \quad E\phi(\underline{x}, \underline{y}) \stackrel{\text{def}}{=} \sum_j p_j \phi(x_j, y_j).$$

Is de verdeling continu, met verdelingsdichtheid $f(x, y)$ dan luidt de definitie

$$(9.13) \quad E\phi(\underline{x}, \underline{y}) \stackrel{\text{def}}{=} \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \phi(x, y) f(x, y) dx dy.$$

DEFINITIE 9.5. De *voorwaardelijke verwachting* van $\phi(\underline{x}, \underline{y})$, onder de voorwaarde $\underline{y} = y$, is de verwachting van $\phi(\underline{x}, y)$ met betrekking tot de voorwaardelijke verdeling van $\underline{x}$ onder dezelfde voorwaarde. *Notatie:* $E(\phi(\underline{x}, \underline{y}) | \underline{y} = y)$ of $E(\phi(\underline{x}, \underline{y}) | y)$, of kortweg: $E(\phi | y)$. Dus, volgens (7.38), voor het continue geval:

$$(9.14) \quad E(\phi | y) = \int_{-\infty}^{+\infty} \phi(x, y) \frac{f(x, y)}{f_2(y)} dx.$$

VOORBEELD. Volgens (7.40) is de voorwaardelijke verdeling van $\underline{x}$, onder de voorwaarde $\underline{y} = y$, als $\underline{x}$ en $\underline{y}$ een simultane normale verdeling (7.28) bezitten, ook normaal. Uit (7.40) en (9.9) volgt nu direct:

$$(9.15) \quad E(\underline{x} | y) = \mu_1 + \rho \frac{\sigma_1}{\sigma_2} (y - \mu_2).$$

DEFINITIE 9.6. De *marginale verwachting* van $\phi(\underline{x}, y)$ met betrekking tot $\underline{x}$, is de verwachting van $\phi(\underline{x}, y)$ genomen over de marginale verdeling van $\underline{x}$. *Notatie:* $E_{\underline{x}} \phi(\underline{x}, y)$.

OPMERKING. $E(\phi | y)$ is een functie van y , zeg $g(y)$; substitueert men hierin $\underline{y}$ voor y , dan verkrijgt men een nieuwe stochastische grootheid $g(\underline{y})$, waarvoor ook de notatie $E(\phi | \underline{y})$ wordt gebruikt. $E_{\underline{x}} \phi(\underline{x}, y)$ is ook een functie van y . Na substitutie van $\underline{y}$ voor y is $E_{\underline{x}} \phi(\underline{x}, \underline{y})$ eveneens een stochastische grootheid, die echter in het algemeen verschilt van de eerstgenoemde; zij zijn gelijk als $\underline{x}$ en $\underline{y}$ stochastisch onafhankelijk zijn (zoals de lezer zelf kan nagaan).

STELLING 9.1. *Hebben $\underline{x}$ en $\underline{y}$ een simultane verdeling en is $\phi(x, y)$ een functie van x alleen, zeg $\phi(x, y) = \psi(x)$, dan is de verwachting $E\phi(\underline{x}, \underline{y})$ van $\phi(\underline{x}, \underline{y})$ gelijk aan de marginale verwachting $E_{\underline{x}} \phi(\underline{x}, \underline{y})$*

BEWIJS. Wij geven het bewijs alleen voor het continue geval; voor het discrete verloopt het analoog.

Volgens (9.13) is

$$\begin{aligned}
E\phi(\underline{x}, \underline{y}) &= E\psi(\underline{x}) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \psi(\underline{x}) f(\underline{x}, \underline{y}) d\underline{x} d\underline{y} = \\
&= \int_{-\infty}^{+\infty} \psi(\underline{x}) d\underline{x} \int_{-\infty}^{+\infty} f(\underline{x}, \underline{y}) d\underline{y} = \int_{-\infty}^{+\infty} \psi(\underline{x}) f_1(\underline{x}) d\underline{x} = \\
&= \int_{-\infty}^{+\infty} \phi(\underline{x}, \underline{y}) f_1(\underline{x}) d\underline{x} = E_{\underline{x}} \phi(\underline{x}, \underline{y}),
\end{aligned}$$

waarin $f_1(\underline{x})$ (vergelijk (7.33)) de marginale verdelingsdichtheid van $\underline{x}$ is. $\square$

VOORBEELD. Voor de tweedimensionale normale verdeling, die gegeven wordt door (7.28), is de marginale verdeling van $\underline{x}$ volgens (7.35) een normale, met parameters μ_1 en σ_1 . Volgens (9.9) en stelling 9.1 is dus voor deze verdeling

$$(9.16) \quad E\underline{x} = \mu_1, \quad E\underline{x}^2 = \sigma_1^2 + \mu_1^2$$

en analoog voor $\underline{y}$.

STELLING 9.2. Zijn $\underline{x}$ en $\underline{y}$ onafhankelijk verdeeld, dan geldt voor de verwachting van een product $\phi(\underline{x})\psi(\underline{y})$ van functies van $\underline{x}$ en $\underline{y}$ afzonderlijk:

$$(9.17) \quad E\phi(\underline{x})\psi(\underline{y}) = E\phi(\underline{x})E\psi(\underline{y}).$$

BEWIJS (voor het continue geval). Met behulp van (8.9) en stelling 9.1 verkrijgen wij:

$$\begin{aligned}
E\phi(\underline{x})\psi(\underline{y}) &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \phi(\underline{x})\psi(\underline{y}) f(\underline{x}, \underline{y}) d\underline{x} d\underline{y} = \\
&= \int_{-\infty}^{+\infty} \phi(\underline{x}) f_1(\underline{x}) d\underline{x} \int_{-\infty}^{+\infty} \psi(\underline{y}) f_2(\underline{y}) d\underline{y} = E\phi(\underline{x})E\psi(\underline{y}). \quad \square
\end{aligned}$$

VOORBEELD. Zijn $\underline{x}$ en $\underline{y}$ onafhankelijk verdeeld, dan is

$$9.18) \quad E\underline{x}\underline{y} = E\underline{x}E\underline{y}.$$

OPMERKING. Het omgekeerde geldt niet, zoals door een tegenvoorbeeld gemakkelijk aan te tonen is. Beschouwen wij bijv. de verdelingsdichtheid (7.29), een homogene verdeling op een cirkelschijf, dan is gemakkelijk in te zien, dat

$$E\bar{x}\bar{y} = E\bar{x} = E\bar{y} = 0$$

is, zodat (9.18) geldt. Dit volgt direct uit de symmetrie van de verdeling ten opzichte van de assen. Toch zijn $\bar{x}$ en $\bar{y}$ stochastisch afhankelijk. (Waarom?)

STELLING 9.3. De operator E is een lineaire operator, d.w.z.

$$(9.19) \quad E\{a\phi(\bar{x}, \bar{y}) + b\psi(\bar{x}, \bar{y}) + c\} = aE\phi(\bar{x}, \bar{y}) + bE\psi(\bar{x}, \bar{y}) + c,$$

als $\bar{x}$ en $\bar{y}$ een simultane verdeling bezitten, ϕ en ψ functies en a , b c constanten zijn.

De stelling geldt analoog voor meer dan 2 stochastische grootheden.

BEWIJS. De geldigheid van (9.19) volgt onmiddellijk uit definitie 9.1. $\square$

TOEPASSING. Deze stelling maakt het mogelijk de gereduceerde momenten in de gewone uit te drukken. Bijvoorbeeld (vergelijk (9.11)):

$$\begin{aligned} \sigma^2(\bar{x}) &= E(\bar{x} - \mu)^2 = E(\bar{x}^2 - 2\mu\bar{x} + \mu^2) = \\ &= E\bar{x}^2 - 2\mu E\bar{x} + \mu^2 = \mu_2 - 2\mu^2 + \mu^2 = \mu_2 - \mu^2, \end{aligned}$$

hetgeen dus leidt tot de volgende zeer nuttige formule ter berekening van de variantie

$$(9.20) \quad \sigma^2 = \mu_2 - \mu^2 \quad \text{of} \quad \sigma^2(\bar{x}) = E\bar{x}^2 - (E\bar{x})^2.$$

OPMERKING. Uit (9.20) volgt direct de ongelijkheid

$$(9.21) \quad \mu_2 \geq \mu^2 \quad \text{of} \quad E\bar{x}^2 \geq (E\bar{x})^2,$$

waarbij de gelijkheidstekens alleen gelden als $\underline{x}$ met kans 1 een bepaalde waarde aanneemt.

STELLING 9.4 (Optelregel voor varianties)

Zijn $\underline{x}$ en $\underline{y}$ stochastisch onafhankelijk verdeeld^{*)} dan geldt

$$(9.22) \quad \sigma^2(\underline{x} + \underline{y}) = \sigma^2(\underline{x}) + \sigma^2(\underline{y}).$$

De stelling geldt ook voor meer dan twee grootheden.

BEWIJS. Volgens (9.19) is

$$E(\underline{x} + \underline{y}) = E\underline{x} + E\underline{y},$$

dus

$$\{E(\underline{x} + \underline{y})\}^2 = (E\underline{x})^2 + (E\underline{y})^2 + 2E\underline{x}E\underline{y}.$$

Verder is, eveneens volgens (9.19)

$$E(\underline{x} + \underline{y})^2 = E(\underline{x}^2 + \underline{y}^2 + 2\underline{x}\underline{y}) = E\underline{x}^2 + E\underline{y}^2 + 2E\underline{x}E\underline{y},$$

daar wegens de onafhankelijkheid (9.18) toegepast kan worden. Volgens (9.20) is

$$\sigma^2(\underline{x} + \underline{y}) = E(\underline{x} + \underline{y})^2 - \{E(\underline{x} + \underline{y})\}^2$$

en substitueren wij de bovenstaande resultaten hierin, dan komt er

$$\sigma^2(\underline{x} + \underline{y}) = E\underline{x}^2 - (E\underline{x})^2 + E\underline{y}^2 - (E\underline{y})^2 = \sigma^2(\underline{x}) + \sigma^2(\underline{y})$$

weer volgens (9.20). $\square$

OPMERKING. De voorwaarde van onafhankelijkheid is voldoende, doch niet noodzakelijk. Nodig en voldoende is blijkens het bewijs, dat (9.18) geldt. Deze stelling geeft tevens een belangrijk verschil aan tussen verwachtingen en varianties. Voor optelbaarheid van verwachtingen behoeft (9.18) niet te gelden.

^{*)} Deze uitdrukking houdt stilzwijgend in, dat $\underline{x}$ en $\underline{y}$ een simultane verdeling bezitten.

TOEPASSING. De stellingen 9.3 en 9.4 stellen ons in staat op eenvoudige wijze de verwachting en de variantie van de binomiale en de negatief-binomiale verdeling te berekenen evenals de verwachting van de hypergeometrische.

Volgens blz. 53 onderaan kan een *binomiaal verdeelde* grootheid $\underline{x}$ (vergelijk (7.5)) beschouwd worden als de som van n onafhankelijk verdeelde grootheden $\underline{x}_i$, die ieder een kans p resp. q bezitten op de waarden 1 resp. 0.

Volgens (9.5) geldt voor ieder daarvan

$$E\underline{x}_i = p \quad \text{en} \quad E\underline{x}_i^2 = p,$$

dus (vergelijk (9.20))

$$\sigma^2(\underline{x}_i) = p - p^2 = pq.$$

Voor

$$\underline{x} = \sum_{i=1}^n \underline{x}_i$$

geldt dus volgens stelling 9.3

$$(9.23) \quad \mu = E\underline{x} = np$$

en volgens (9.22)

$$(9.24) \quad \sigma^2 = \sigma^2(\underline{x}) = npq. *)$$

Voor de *negatief-binomiale* verdeling (formule (7.7)) verloopt de berekening analoog. Stelt $\underline{x}$ het aantal experimenten voor tot en met het k -de succes, dan kan $\underline{x}$ beschouwd worden als de som van k onafhankelijke grootheden $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_k$ waarvan $\underline{x}_1$ het aantal experimenten voorstelt tot en met het eerste succes, $\underline{x}_2$ het daarop volgende aantal experimenten tot en met het tweede succes, enz. Daar alle experimenten onafhankelijk zijn, geldt dit ook voor de grootheden $\underline{x}_i$ ($i = 1, 2, \dots, k$) en deze bezitten bovendien alle

*) Deze variantie wordt veelal berekend door $E\underline{x}(\underline{x}-1)$ direct uit de definiërende formule $E\underline{x}(\underline{x}-1) = \sum_{x=0}^n x(x-1)p^x q^{n-x} \binom{n}{x}$ te berekenen en vervolgens hieruit $E\underline{x}^2$ hieruit af te leiden. Deze berekeningswijze is ingewikkelder dan de hier gevolgde.

dezelfde verdeling als $\underline{x}_1$. De momenten daarvan worden gegeven door (9.6) en (9.7), dus

$$E\underline{x}_i = \frac{1}{p} \quad \text{en} \quad E\underline{x}_i^2 = \frac{1+q}{p^2},$$

waaruit volgens (9.20) volgt

$$\sigma^2(\underline{x}_i) = \frac{1+q}{p^2} - \frac{1}{p^2} = \frac{q}{p^2}.$$

Derhalve leiden stelling 9.3 en 9.4 tot

$$(9.25) \quad \mu = E\underline{x} = \frac{k}{p} \quad \text{en} \quad \sigma^2 = \sigma^2(\underline{x}) = \frac{kq}{p^2}.$$

Voor de *hypergeometrische verdeling* (vergelijk (7.6)) kan deze redenering slechts gedeeltelijk gevolgd worden. Volgens voorbeeld 7 van hoofdstuk V wordt deze verdeling verkregen als die van het aantal witte knikkers bij aselechte trekking zonder teruglegging van n knikkers uit een doos met M witte en N zwarte knikkers. Noemen wij nu het aantal witte knikkers bij de i -de trekking $\underline{x}_i$ (met mogelijke waarden 0 en 1), dan geldt weer

$$\underline{x} = \sum_{i=1}^n \underline{x}_i$$

en volgens voorbeeld 5 (blz. 40) bezitten alle $\underline{x}_i$ dezelfde (marginale) verdeling, met een kans $\frac{M}{M+N}$ op de waarde 1. Zij zijn nu echter niet stochastisch onafhankelijk (waarom niet?), zodat (9.22) niet toegepast kan worden. Wel echter stelling 9.3 en deze geeft, daar $E\underline{x}_i = \frac{M}{M+N}$ is

$$(9.26) \quad \mu = E\underline{x} = \frac{Mn}{M+N}.$$

Voor de *normale verdeling* volgt uit (9.9) en (9.22)

$$(9.27) \quad \sigma^2(\underline{x}) = \sigma^2 + n^2 - \mu^2 = \sigma^2.$$

Verwachting en variantie van de normale verdeling zijn dus juist de parameters μ en σ^2 , die in formule (7.13) voorkomen. Een normale verdeling wordt dus door verwachting en variantie volledig bepaald.

STELLING 9.5. Zijn a en b constanten en is $\underline{x}$ een stochastische grootte,

dan is

$$(9.28) \quad \sigma^2(a\underline{x}+b) = a^2 \sigma^2(\underline{x}).$$

BEWIJS. $E(a\underline{x}+b) = aE\underline{x} + b.$

De bij $a\underline{x} + b$ behorende gereduceerde grootheid is dus

$$a\underline{x} + b - aE\underline{x} - b = a(\underline{x} - E\underline{x}) = a\check{\underline{x}}.$$

Dus

$$\sigma^2(a\underline{x}+b) = E(a\check{\underline{x}})^2 = Ea^2\check{\underline{x}}^2 = a^2E\check{\underline{x}}^2 = a^2\sigma^2(\underline{x}). \quad \square$$

De belangrijkste toepassing hiervan is vervat in:

STELLING 9.6. *Is $\underline{x}$ een stochastische grootheid met verwachting μ en variantie σ^2 , dan bezit de stochastische grootheid*

$$(9.29) \quad \underline{x}^* \stackrel{\text{def}}{=} \frac{\underline{x} - \mu}{\sigma} = \frac{\check{\underline{x}}}{\sigma}$$

verwachting 0 en variantie 1.

Het bewijs wordt aan de lezer overgelaten.

OPMERKINGEN. De grootheid $\underline{x}^*$ wordt de bij $\underline{x}$ behorende *gestandaardiseerde grootheid* of kortweg de *gestandaardiseerde $\underline{x}$* genoemd (vergelijk ook blz. 51).

Dit standaardiseringsproces stelt ons o.a. in staat om iedere normale verdeling te beheersen met behulp van een tabel van de gestandaardiseerde normale verdeling.

VOORBEELD. $\underline{x}$ is $N(4;0,5)$ -verdeeld; dus normaal met verwachting 4 en spreiding 0,5. Gevraagd wordt te berekenen

$$P(\underline{x} \geq 4,75).$$

De berekening verloopt als volgt:

$$P(\underline{x} \geq 4,75) = P\left(\frac{\underline{x} - \mu}{\sigma} \geq \frac{4,75 - \mu}{\sigma}\right) = P(\underline{u} \geq 1,5),$$

waarin $\underline{u}$ de gestandaardiseerde normaal verdeelde grootheid voorstelt. Uit de tabel van de normale verdeling (tabel 1) blijkt, dat de gezochte kans 0,067 bedraagt. Verdere opgaven van deze aard zijn in het opgavenboekje (zie voetnoot op blz. 67) te vinden.

STELLING 9.7. Indien $\underline{x}$ en $\underline{y}$ een simultane verdeling bezitten en $\phi(\underline{x}, \underline{y})$ is een functie van $\underline{x}$ en $\underline{y}$, terwijl

$$(9.30) \quad g(\underline{y}) \stackrel{\text{def}}{=} E(\phi(\underline{x}, \underline{y}) | \underline{y} = \underline{y})$$

de voorwaardelijke verwachting van $\phi(\underline{x}, \underline{y})$ onder de voorwaarde $\underline{y} = \underline{y}$ is, dan geldt

$$(9.31) \quad E\phi(\underline{x}, \underline{y}) = Eg(\underline{y}).$$

BEWIJS (voor het continue geval).

Volgens definitie 9.4 en 9.5 en stelling 9.1 geldt

$$\begin{aligned} E\phi(\underline{x}, \underline{y}) &= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \phi(\underline{x}, \underline{y}) f(\underline{x}, \underline{y}) d\underline{x} d\underline{y} = \\ &= \int_{-\infty}^{+\infty} f_2(\underline{y}) d\underline{y} \int_{-\infty}^{+\infty} \phi(\underline{x}, \underline{y}) \frac{f(\underline{x}, \underline{y})}{f_2(\underline{y})} d\underline{x} = \\ &= \int_{-\infty}^{+\infty} g(\underline{y}) f_2(\underline{y}) d\underline{y} = Eg(\underline{y}). \quad \square \end{aligned}$$

TOEPASSINGEN. Met behulp van deze stelling kunnen wij nu de variantie van een hypergeometrische verdeling berekenen.*¹ Laat daartoe $\underline{x}_i$ resp. $\underline{x}_j$ de aantallen witte knikkers bij de i -de resp. j -de trekking zonder teruglegging voorstellen (beide nemen dus één van de waarden 0 of 1 aan), dan berekenen wij eerst $E\underline{x}_i \underline{x}_j$. Op analoge wijze als bij voorbeeld 5 van hoofdstuk V valt in te zien, dat de verdeling van $\underline{x}_i$, bij gegeven $\underline{x}_j$, voor alle $i \neq j$ dezelfde is en wel

*¹) Ook deze variantie kan, evenals die van de binomiale verdeling, berekend worden door de definitie 9.1 toe te passen op $E\underline{x}(\underline{x}-1)$.

$$P(\underline{x}_i = 1 | \underline{x}_j = 1) = \frac{M-1}{M+N-1} ,$$

en

$$P(\underline{x}_i = 1 | \underline{x}_j = 0) = \frac{M}{M+N-1} .$$

Derhalve geldt:

$$E(\underline{x}_i | \underline{x}_j = 1) = \frac{M-1}{M+N-1} ; \quad E(\underline{x}_i | \underline{x}_j = 0) = \frac{M}{M+N-1} .$$

Dus is

$$E(\underline{x}_i \underline{x}_j | \underline{x}_j = 1) = E(\underline{x}_i | \underline{x}_j = 1) = \frac{M-1}{M+N-1}$$

en

$$E(\underline{x}_i \underline{x}_j | \underline{x}_j = 0) = E(0 | \underline{x}_j = 0) = 0 .$$

Volgens stelling 9.7 is dus

$$(9.32) \quad E \underline{x}_i \underline{x}_j = \frac{M-1}{M+N-1} P(\underline{x}_j = 1) + 0 = \frac{M-1}{M+N-1} \cdot \frac{M}{M+N} ; \quad *)$$

Voor $\underline{x} = \sum_{i=1}^n \underline{x}_i$ geldt nu

$$E \underline{x}^2 = E \left(\sum_{i=1}^n \underline{x}_i \right)^2 = E \sum_{i=1}^n \underline{x}_i^2 + 2E \sum_{i < j} \underline{x}_i \underline{x}_j ,$$

maar hierin is volgens (9.5) voor iedere i

$$E \underline{x}_i = E \underline{x}_i^2 = \frac{M}{M+N}$$

en volgens (9.32) voor ieder paar i, j met $i \neq j$

$$E \underline{x}_i \underline{x}_j = \frac{M(M-1)}{(M+N)(M+N-1)} .$$

Dus is

*) (9.32) kan natuurlijk rechtstreeks berekend worden, namelijk:

$$E \underline{x}_i \underline{x}_j = P(\underline{x}_i = 1 \wedge \underline{x}_j = 1) =$$

$$P(\underline{x}_i = 1 | \underline{x}_j = 1) P(\underline{x}_j = 1) = \frac{M-1}{M+N-1} \cdot \frac{M}{M+N} ;$$

$$E\bar{x}^2 = \frac{nM}{M+N} + \frac{n(n-1)M(M-1)}{(M+N)(M+N-1)}.$$

Toepassing van (9.20), tezamen met (9.26) geeft dan na enige herleiding

$$(9.33) \quad \sigma^2(\bar{x}) = \frac{nMN(M+N-n)}{(M+N)^2(M+N-1)}.$$

Een grootheid, die evenals verwachting en variantie van groot belang is, is de *covariantie* van twee stochastische grootheden $\bar{x}$ en $\bar{y}$. Deze wordt geschreven als $\text{cov}(\bar{x}, \bar{y})$ en gedefinieerd door

$$(9.34) \quad \text{cov}(\bar{x}, \bar{y}) \stackrel{\text{def}}{=} E\bar{x}\bar{y} = E(\bar{x} - E\bar{x})(\bar{y} - E\bar{y}).$$

STELLING 9.8. Zijn $\bar{x}$ en $\bar{y}$ onafhankelijk verdeeld, dan is

$$(9.35) \quad \text{cov}(\bar{x}, \bar{y}) = 0.$$

BEWIJS. Men houde in het oog, dat $E\bar{x}$ en $E\bar{y}$ constanten zijn. Volgens stelling 9.3 is

$$\begin{aligned} E(\bar{x} - E\bar{x})(\bar{y} - E\bar{y}) &= E(\bar{x}\bar{y} - \bar{x}E\bar{y} - \bar{y}E\bar{x} + E\bar{x}E\bar{y}) = \\ &= E\bar{x}\bar{y} - E\bar{x}E\bar{y} - E\bar{y}E\bar{x} + E\bar{x}E\bar{y} = E\bar{x}\bar{y} - E\bar{x}E\bar{y} = 0 \end{aligned}$$

volgens (9.18).

Wij berekenen nu, met behulp van stelling 9.7 de covariantie van $\bar{x}$ en $\bar{y}$, voor het geval van de *tweedimensionale normale verdeling* (formule (7.28)). Volgens (9.15) is

$$E(\bar{x} - \mu_1 | \bar{y}) = \rho \frac{\sigma_1}{\sigma_2} (\bar{y} - \mu_2),$$

dus (vgl. (9.30))

$$g(\bar{y}) = E\{(\bar{x} - \mu_1)(\bar{y} - \mu_2) | \bar{y}\} = \rho \frac{\sigma_1}{\sigma_2} (\bar{y} - \mu_2)^2.$$

Volgens (9.31) is dus

$$\begin{aligned} E(\bar{x} - \mu_1)(\bar{y} - \mu_2) &= E\rho \frac{\sigma_1}{\sigma_2} (\bar{y} - \mu_2)^2 = \rho \frac{\sigma_1}{\sigma_2} E(\bar{y} - \mu_2)^2 = \\ &= \rho \frac{\sigma_1}{\sigma_2} \sigma_2^2 = \rho \sigma_1 \sigma_2. \end{aligned}$$

Dus

$$(9.36) \quad \text{cov}(\underline{x}, \underline{y}) = \rho \sigma_1 \sigma_2.$$

STELLING 9.9. *Bezitten $\underline{x}$ en $\underline{y}$ een simultane verdeling, dan geldt*

$$(9.37) \quad \sigma^2(a\underline{x} + b\underline{y}) = a^2 \sigma^2(\underline{x}) + b^2 \sigma^2(\underline{y}) + 2ab \text{cov}(\underline{x}, \underline{y}).$$

BEWIJS. De bij $a\underline{x} + b\underline{y}$ behorende gereduceerde grootheid is $a\underline{\tilde{x}} + b\underline{\tilde{y}}$, daar $E(a\underline{x} + b\underline{y}) = aE\underline{x} + bE\underline{y}$ is. Derhalve is

$$\sigma^2(a\underline{x} + b\underline{y}) = E(a\underline{\tilde{x}} + b\underline{\tilde{y}})^2 = a^2 E\underline{\tilde{x}}^2 + b^2 E\underline{\tilde{y}}^2 + 2ab E\underline{\tilde{x}}\underline{\tilde{y}}. \quad \square$$

STELLING 9.10. *Bezitten $\underline{x}$ en $\underline{y}$ een simultane verdeling, dan geldt*

$$(9.38) \quad |\text{cov}(\underline{x}, \underline{y})| \leq \sigma(\underline{x})\sigma(\underline{y}).$$

BEWIJS. Stel in (9.37) $b = 1$. Daar een variantie steeds ≥ 0 is, geldt

$$0 \leq \sigma^2(a\underline{x} + \underline{y}) = \sigma^2(\underline{x})a^2 + 2\text{cov}(\underline{x}, \underline{y})a + \sigma^2(\underline{y}).$$

Dit geldt voor iedere a , dus moet de discriminant van deze kwadratische vorm in a niet positief zijn:

$$\frac{1}{4}D = \{\text{cov}(\underline{x}, \underline{y})\}^2 - \sigma^2(\underline{x})\sigma^2(\underline{y}) \leq 0.$$

Hieruit volgt (9.38). $\square$

De *correlatiecoëfficiënt* van $\underline{x}$ en $\underline{y}$ wordt gedefinieerd door

$$(9.39) \quad \rho(\underline{x}, \underline{y}) \stackrel{\text{def}}{=} \frac{\text{cov}(\underline{x}, \underline{y})}{\sigma(\underline{x})\sigma(\underline{y})}.$$

Volgens (9.38) geldt

$$(9.40) \quad -1 \leq \rho(\underline{x}, \underline{y}) \leq +1$$

en volgens (9.36) is $\rho(\underline{x}, \underline{y})$ voor een tweedimensionale normale verdeling gelijk aan de parameter ρ uit formule (7.28).

OPGAVE. Twee stochastische grootheden $\underline{x}$ en $\underline{y}$ bezitten een simultane kansverdeling met $|\rho(\underline{x}, \underline{y})| = 1$. Dit betekent, dat het stochastische punt $(\underline{x}, \underline{y})$ met kans 1 op een bepaalde rechte lijn ligt, en omgekeerd. Bewijs dit. (Dit is opgave 88 van het opgavenboekje. Zie voor de oplossing blz. 131-132 daarvan.)

HOOFDSTUK X

DE ONGELIJKHEID VAN BIENAYMÉ-TCHEBYCHEFF EN DE THEORETISCHE WET VAN DE GROTE AANTALLEN

Een vooral voor theoretische doeleinden belangrijke ongelijkheid is door BIENAYMÉ en door TCHEBYCHEFF (onafhankelijk van elkaar) gevonden.

STELLING 10.1. Voor een stochastische grootheid $\underline{x}$ met verwachting μ en spreiding σ geldt voor iedere $k > 0$

$$(10.1) \quad P(|\underline{x}^*| \geq k) \leq \frac{1}{k^2},$$

waarin

$$(10.2) \quad \underline{x}^* = \frac{\underline{x} - \mu}{\sigma}$$

de gestandaardiseerde $\underline{x}$ voorstelt.

BEWIJS. Wij geven het bewijs voor een discreet verdeelde $\underline{x}$; voor een continue verdeling verloopt het analoog. Per definitie is

$$\sigma^2 = \sum_i p_i (x_i - \mu)^2,$$

als x_i ($i = 1, 2, \dots$) de waarden voorstellen, die $\underline{x}$ aan kan nemen. Dus

$$1 = \sum_i p_i \frac{(x_i - \mu)^2}{\sigma^2} = \sum_i p_i x_i^{*2}.$$

Deze som splitsen wij in twee deelsommen Σ_1 en Σ_2 , waarbij

$$\begin{array}{ll} \Sigma_1 & \text{loopt over die } i, \text{ waarvoor } |x_i^*| \geq k \\ \text{en } \Sigma_2 & \text{" " " } i, \text{ " } |x_i^*| < k. \end{array}$$

Dan is uiteraard

$$1 \geq \sum_1 p_i x_i^{*2} \geq k^2 \sum_1 p_i.$$

Hierin is echter

$$\sum_1 p_i = P(|\underline{x}^*| \geq k),$$

waaruit (10.1) direct volgt. $\square$

VOORBEELDEN. Laat $\underline{x}$ de waarden 1 en 0 met kansen p en $q = 1 - p$ aannemen, zodat $\mu = p$ en $\sigma^2 = pq$ is. Dan geldt volgens (10.1)

$$P(|\underline{x} - p| \geq k\sqrt{pq}) \leq \frac{1}{k^2}.$$

Vullen wij $p = \frac{1}{2}$ en $k = 1$ in, dan komt er

$$P(|\underline{x} - \frac{1}{2}| \geq \frac{1}{2}) \leq 1,$$

hetgeen in dit geval in een gelijkheid overgaat, daar de kans in het linkerlid = 1 is. In dit speciale geval is de ongelijkheid dus scherp: verkleining van het rechterlid is niet mogelijk zonder tot een onjuist resultaat te komen.

In de regel is de ongelijkheid echter zeer grof. Bezit $\underline{x}$ bijv. een $N(0,1)$ -verdeling, dan geldt volgens (10.1)

$$P(|\underline{x}| \geq 2) \leq \frac{1}{4},$$

terwijl volgens tabel 1 de kans in het linkerlid gelijk is aan 0,046.

Voor theoretische doeleinden is dit gebrek aan scherpheid veelal van weinig belang, zoals uit de hieronder vermelde stellingen blijkt.

STELLING 10.2. Als $\underline{x}_1, \underline{x}_2, \dots$ onafhankelijk verdeelde stochastische grootheden zijn, die alle dezelfde verwachting μ en variantie σ^2 bezitten, dan convergeert, voor $n \rightarrow \infty$, het gemiddelde

$$(10.3) \quad \underline{\bar{x}}_n \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n \underline{x}_i$$

stochastisch naar μ , d.w.z. voor iedere $\epsilon > 0$ geldt

$$(10.4) \quad \lim_{n \rightarrow \infty} P(|\bar{x}_n - \mu| > \epsilon) = 0. \quad *)$$

BEWIJS. Volgens stelling 9.3 is

$$E \sum_{i=1}^n x_i = n\mu,$$

dus

$$(10.5) \quad E\bar{x}_n = \mu.$$

Volgens stelling 9.4 is

$$\sigma^2 \left(\sum_{i=1}^n x_i \right) = n\sigma^2$$

en dus volgens stelling 9.5

$$(10.6) \quad \sigma^2(\bar{x}_n) = \frac{\sigma^2}{n}.$$

Invullen in (10.1) met $k = \frac{\sqrt{n}\epsilon}{\sigma}$ geeft

$$P\left(\left|\frac{\bar{x} - \mu}{\sigma}\right| \geq \frac{\sqrt{n}\epsilon}{\sigma}\right) = P(|\bar{x}_n - \mu| \geq \epsilon) \leq \frac{\sigma^2}{n\epsilon^2}$$

en de limiet, voor $n \rightarrow \infty$, van het rechterlid is 0. $\square$

Een speciaal geval van deze stelling wordt verkregen door voor de x_i het aantal successen (0 of 1) te nemen bij een experiment met kans p op succes. De stelling neemt dan de volgende vorm aan.

STELLING 10.3 (Theoretische wet der grote aantallen)

In een reeks van n onafhankelijke experimenten, ieder met kans p op succes, convergeert de fractie successen $\bar{x}/n$ voor $n \rightarrow \infty$ in waarschijnlijkheid naar p , d.w.z. voor iedere $\epsilon > 0$ is

$$(10.7) \quad \lim_{n \rightarrow \infty} P\left(\left|\frac{\bar{x}}{n} - p\right| \geq \epsilon\right) = 0.$$

*) Men noemt dit ook wel: convergentie in waarschijnlijkheid.

OPMERKINGEN. Deze stelling heeft een grote praktische betekenis. Het is de theoretische tegenhanger van de experimentele wet van de grote aantallen (zie hoofdstuk II). Het zou niet juist zijn te zeggen, dat deze laatste nu "bewezen" is, daar een stelling, die op de werkelijkheid betrekking heeft, nooit op grond van een axiomastelsel bewezen kan worden. Alleen de experimentele bevestiging is daar geschikt voor. Wel kunnen wij nu echter opmerken, dat de uitermate eenvoudige axioma's, waarvan wij zijn uitgegaan, adaequaat blijken te zijn ter beschrijving van de nogal ingewikkelde experimentele wet van de grote aantallen, die wij in hoofdstuk II fundamenteel voor statistische verschijnselen hebben genoemd. Dit betekent, dat de interpretatie van een kans p als - bij benadering - een f_q in een lange reeks waarnemingen bevestigd is, terwijl bovendien dit succes van de theorie het vertrouwen in verdere resultaten - d.w.z. het vertrouwen dat de toepassingen daarvan in de praktijk bruikbaar zullen blijken te zijn - vergroot. De op blz. 28 gemaakte opmerking, "dat in het kansmodel het f_q als zodanig, met zijn statistische fluctuaties, op natuurlijke wijze weer ingevoerd kan worden", is hiermede waar gemaakt.

De praktische betekenis van stelling 10.2 is de volgende. Indien men van een stochastische grootheid $\underline{x}$ met verwachting μ en variantie σ^2 , een aantal onafhankelijke waarnemingen verricht, kan de i -de waarneming voorgesteld worden door $\underline{x}_i$ en is de stelling toepasbaar, daar alle $\underline{x}_i$ nu dezelfde verdeling (n.l. dezelfde verdeling als $\underline{x}$) bezitten. Formule (10.3) betekent nu, dat het rekenkundige gemiddelde der waarnemingen, naarmate n groter wordt, steeds minder kans bezit om ver van μ verwijderd te zijn. Als men n maar groot genoeg maakt, kan men - behoudens een willekeurig kleine kans - zorgen willekeurig dicht bij μ terecht te komen.

Indien men bijv. een fysische of chemische constante (gewicht, soortelijk gewicht, titer van een vloeistof, lading van een electron, enz.) nauwkeurig wenst te bepalen, kan men een aantal onafhankelijke waarnemingen daarvan verrichten. Ieder daarvan is aan meetfouten onderhevig en kan daarom als een stochastische grootheid beschouwd worden. Worden de waarnemingen op dezelfde wijze (door dezelfde personen en met dezelfde apparatuur) verricht, dan bezitten ze alle dezelfde kansverdeling. Worden bovendien de uitkomsten van latere waarnemingen niet beïnvloed door die van de voorafgaande - gevaar voor psychologische beïnvloeding is vaak in sterke mate aanwezig -, dan kunnen zij als stochastisch onafhankelijk beschouwd worden. In deze omstandigheden is de stelling toepasbaar. Veelal zal men

n dan niet zeer groot kunnen (of willen) maken, maar ook bij kleine n doet de invloed van formule (10.6) zich reeds voor: de variabiliteit van een gemiddelde van een aantal onafhankelijke waarnemingen is aanzienlijk kleiner, dus de nauwkeurigheid aanzienlijk groter dan van één enkele waarneming.

De grootheid μ , waarbij $\bar{x}_n$ in de buurt komt, hangt vaak mede van de waarnemingsapparatuur af. Men dient zich daarvan bij het verrichten van metingen steeds bewust te zijn (hetgeen niet altijd het geval is). Bij de bepaling van een titer van een vloeistof bijv. kunnen afwijkingen in de gebruikte buret en pipet systematische verschillen veroorzaken tussen μ en de titer, die men eigenlijk wenst te meten. Het is dan zinloos de nauwkeurigheid van de bepaling, door n op te voeren, zeer groot te maken, daar de systematische fout dan gaat overheersen. Men kan daaraan tegemoet komen door de apparatuur te wisselen.

Laat men verschillende waarnemers werken met dezelfde of verschillende apparatuur, maar zijn er geen (of nauwelijks) systematische verschillen, zodat μ voor alle waarnemingen gelijk is, dan is het van belang op te merken, dat stelling 10.2 ook blijft gelden als de varianties ongelijk zijn, mits deze begrensd zijn. Is $\sigma^2(\underline{x}_i) \leq \sigma_0^2$ voor alle i , dan is n.l. het bewijs, met een kleine wijziging, aan te passen. De voorwaarden kunnen trouwens nog verder verzwakt worden, doch dit punt zullen wij hier buiten beschouwing laten.

De uitwerking van de hierboven geschetste gedachtengang behoort tot de verderop te behandelen *schattingstheorie*.

HOOFDSTUK XI

FUNCTIES VAN STOCHASTISCHE GROOTHEDEN; CENTRALE LIMIETSTELLING

In de vorige hoofdstukken is reeds opgemerkt, dat een functie van één of meer stochastische grootheden zelf weer een stochastische grootheid is. Eén der belangrijkste voorbeelden hiervan is het gemiddelde van een aantal stochastische grootheden, waarvan in hoofdstuk X sprake is geweest.

Wij gaan nu na, hoe in het algemeen de kansverdeling van een dergelijke functie berekend kan worden.

STELLING 11.1. Laat $x_1, x_2, \dots, x_n$ een simultane verdeling bezitten met $f(x_1, x_2, \dots, x_n)$ als kansdichtheid en laat $\phi(x_1, x_2, \dots, x_n)$ een functie van $x_1, x_2, \dots, x_n$ zijn. Laat verder, in de n -dimensionale ruimte R_n met $x_1, x_2, \dots, x_n$ als coördinaten, G_a het gebied voorstellen van al die punten, waarvoor

$$(11.1) \quad \phi(x_1, x_2, \dots, x_n) \leq a$$

is. Dan is

$$(11.2) \quad \begin{aligned} F_{\phi}(a) &= P(\phi \leq a) = P(\phi(x_1, x_2, \dots, x_n) \leq a) = \\ &= P(\underline{p} \in G_a) = \int_{G_a} \dots \int f(x_1, x_2, \dots, x_n) dx_1 dx_2 \dots dx_n. \end{aligned}$$

Hierin stelt $\underline{p}$ het stochastische punt $(x_1, x_2, \dots, x_n)$ voor. In geval van een discrete verdeling wordt in het laatste lid van (11.2) de meervoudige integraal door een som vervangen. De stelling volgt uit (7.27).

Een toepassing van de stelling is de volgende.

STELLING 11.2. Als x een $N(\mu, \sigma)$ -verdeelde grootheid is, is

$$(11.3) \quad \underline{u} = \frac{\underline{x} - \mu}{\sigma}$$

$N(0,1)$ -verdeeld.

BEWIJS. Het gebied G_a , waarin $u \leq a$ is, wordt gegeven door $\frac{x-\mu}{\sigma} \leq a$ of $x \leq \mu + a\sigma$. Of, direct opgeschreven

$$F_{\underline{u}}(u) \stackrel{\text{def}}{=} P(\underline{u} \leq u) = P(\underline{x} \leq \mu + u\sigma) = F_{\underline{x}}(\mu + u\sigma).$$

De verdelingsdichtheid $f_{\underline{u}}(u)$ wordt hieruit verkregen door differentiëren:

$$f_{\underline{u}}(u) = \frac{d}{du} F_{\underline{u}}(u) = \frac{d}{du} F_{\underline{x}}(\mu + u\sigma) = \sigma f_{\underline{x}}(\mu + u\sigma).$$

Vullen wij echter in $f_{\underline{x}}(x)$ voor x in: $\mu + u\sigma$, dan komt er weer een kwadratische e-macht, dus weer een normale verdelingsdichtheid. Volgens stelling 9.6 heeft $\underline{u}$ bovendien verwachting 0 en spreiding 1, waarmee het bewijs geleverd is. $\square$

Het onder stelling 9.6 gegeven voorbeeld is een toepassing van stelling 11.2, die daar reeds (zonder vermelding) is gebruikt.

Numerieke voorbeelden van toepassing van stelling 11.1 zijn vervat in de volgende twee opgaven, die betrekking hebben op discrete verdelingen.

OPGAVE 11.1. Twee stochastische grootheden $\underline{x}$ en $\underline{y}$ zijn beide binomiaal verdeeld; $\underline{x}$ met parameters $n=3$ en $p=\frac{1}{3}$, $\underline{y}$ met parameters $m=4$ en $p_1=\frac{1}{4}$. Verder is gegeven, dat $\underline{x}$ en $\underline{y}$ onafhankelijk verdeeld zijn.

Bepaal de kansverdeling van $\underline{z} = \underline{x} + \underline{y}$ en bereken verwachting en spreiding van deze verdeling.

(Opgave 72 van het Opgavenboekje; zie voor de oplossing blz. 109-111 daarvan.)

OPGAVE 11.2. Een tentamen bestaat uit 5 vragen, die goed of fout beantwoord kunnen worden. Voor ieder goed antwoord wordt 1 punt gegeven, voor een fout antwoord 0.

- a) Candidaat A heeft zich zodanig op het tentamen voorbereid, dat hij bij iedere vraag een kans $\frac{2}{3}$ heeft op het geven van het goede antwoord.

Bereken, in de onderstelling van onafhankelijkheid der antwoorden, de kansverdeling van zijn eindcijfer en de verwachting en spreiding daarvan.

- b) Doe hetzelfde voor kandidaat B, die voor de eerste 3 vragen dezelfde kans op een goed antwoord heeft als A, maar voor de laatste twee slechts een kans $\frac{1}{2}$.

(Opgave 73 van het Opgavenboekje; zie voor de oplossing blz. 111-113 daarvan.)

Een belangrijke toepassing van stelling 11.1 is de volgende.

STELLING 11.3. Zijn $\underline{x}$ en $\underline{y}$ continu en onafhankelijk verdeelde stochastische grootheden met verdelingsfunctie $F(x)$ en $G(y)$ en verdelingsdichtheden $f(x)$ en $g(y)$, dan geldt voor verdelingsfunctie $H(z)$ en verdelingsdichtheid $h(z)$ van $\underline{z} = \underline{x} + \underline{y}$

$$(11.4) \quad H(z) = \int_{-\infty}^{+\infty} f(x)G(z-x)dx$$

en

$$(11.5) \quad h(z) = \int_{-\infty}^{+\infty} f(x)g(z-x)dx.$$

BEWIJS. Volgens (11.2) en stelling 8.3 is

$$H(z) = \iint_{x+y \leq z} f(x)g(y)dxdy = \int_{-\infty}^{+\infty} f(x)dx \int_{-\infty}^{z-x} g(y)dy = \int_{-\infty}^{+\infty} f(x)G(z-x)dx.$$

Hieruit wordt $h(z)$ verkregen door naar z te differentiëren, hetgeen onder de integraal mag geschieden. Daarbij is dan echter x als een constante op te vatten, zodat (11.5) direct volgt. $\square$

Een toepassing hiervan is de volgende.

STELLING 11.4. Zijn $\underline{x}$ en $\underline{y}$ onafhankelijk normaal verdeeld, dan is $\underline{z} = \underline{x} + \underline{y}$ eveneens normaal verdeeld.

BEWIJS. De verdelingsdichtheden $f(x)$ en $g(y)$ zijn nu van de vorm

$$f(x) :: e^{-\frac{1}{2}\left(\frac{x-\mu_1}{\sigma_1}\right)^2} \quad \text{en} \quad g(y) :: e^{-\frac{1}{2}\left(\frac{y-\mu_2}{\sigma_2}\right)^2},$$

waarin :: betekent: op een constante factor na gelijk aan.

Het product $f(x)g(z-x)$ heeft dus de vorm

$$f(x)g(z-x) :: e^{-\frac{1}{2}\left\{\left(\frac{x}{\sigma_1}\right)^2 - \frac{2\mu_1 x}{\sigma_1^2} + \left(\frac{z-x}{\sigma_2}\right)^2 - \frac{2\mu_2(z-x)}{\sigma_2^2}\right\}}.$$

De exponent kan (vergelijk blz. 57) gesplitst worden in twee delen, waarvan het ene een kwadratische vorm in z is (zonder x) en het tweede een volkomen kwadraat van de vorm $(\alpha x + \beta z + \gamma)^2$. Het eerste deel kan als een e -macht voor de integraal van (11.5) gebracht worden en deze integraal zelf is een constante (substitueer daarin n.l. $v = \alpha x + \beta z + \gamma$, dan veranderen de grenzen niet; deze blijven $-\infty$ en $+\infty$). Daaruit volgt echter de stelling. $\square$

OPGAVE 11.3.

- Geeft de formule voor $h(z)$ uit stelling 11.4.
- Bereken de correlatiecoëfficiënt van de grootheden $\underline{x}$ en $\underline{z}$ uit de stelling 11.4.
- Geef de formule voor de simultane verdelingsdichtheid van $\underline{x}$ en $\underline{z}$.

OPLOSSING van b). De correlatiecoëfficiënt wordt gedefinieerd door (9.39).

Zijn $\tilde{\underline{x}}$ en $\tilde{\underline{y}}$ de gereduceerden van $\underline{x}$ en $\underline{y}$, dus

$$\tilde{\underline{x}} = \underline{x} - \mu_1 \quad \text{en} \quad \tilde{\underline{y}} = \underline{y} - \mu_2$$

en is $\tilde{\underline{z}}$ de gereduceerde $\underline{z}$, dan is volgens (9.19)

$$\tilde{\underline{z}} = \underline{z} - E\underline{z} = \underline{z} - E(\underline{x} + \underline{y}) = \underline{z} - \mu_1 - \mu_2 = \tilde{\underline{x}} + \tilde{\underline{y}}.$$

Derhalve is (vergelijk blz.)

$$\text{cov}(\underline{x}, \underline{z}) = E\tilde{\underline{x}}\tilde{\underline{z}} = E\tilde{\underline{x}}(\tilde{\underline{x}} + \tilde{\underline{y}}) = E\tilde{\underline{x}}^2 = \sigma_1^2,$$

daar $E\tilde{\underline{x}}\tilde{\underline{y}} = 0$ wegens de onafhankelijkheid van $\underline{x}$ en $\underline{y}$.

Dus is, met gebruik maken van stelling 9.4:

$$\rho(\underline{x}, \underline{z}) = \frac{\text{cov}(\underline{x}, \underline{z})}{\sigma(\underline{x})\sigma(\underline{z})} = \frac{\sigma_1^2}{\sigma_1 \sqrt{\sigma_1^2 + \sigma_2^2}} = \sqrt{\frac{\sigma_1^2}{\sigma_1^2 + \sigma_2^2}}$$

$\underline{x}$ en $\underline{z}$ zijn dus niet onafhankelijk verdeeld.

OPMERKINGEN. Stelling 11.4 geldt ook als deze grootheden niet stochastisch onafhankelijk zijn, doch een simultane normale verdeling bezitten. Het bewijs verloopt op geheel analoge wijze met behulp van de formules (7.28) en (11.2), waarin voor G_a het gebied $x+y \leq z$ genomen wordt. Beide stellingen blijven gelden voor meer dan 2 simultaan normaal verdeelde grootheden. Voor het geval van onafhankelijkheid is dit eenvoudiger te bewijzen door volledige inductie.

Tevens is gemakkelijk in te zien, dat vermenigvuldiging van de stochastische grootheden met constante factoren de normaliteit van de simultane verdeling niet verstoort (vergelijk stelling 11.2), evenmin als het optellen van constante termen, zodat de volgende algemene stelling geldt.

STELLING 11.5. *Bezitten $x_1, x_2, \dots, x_n$ een simultane normale verdeling, dan is ook iedere lineaire combinatie*

$$(11.6) \quad z = a_0 + a_1 x_1 + \dots + a_n x_n$$

normaal verdeeld.

OPGAVE 11.4. Een leverancier van flessen slaolie vermeldt als netto inhoud van zijn flessen: 350 gram olie. Een winkelier wenst deze bewering te controleren zonder de flessen te openen. Hij weegt daartoe een zeer groot aantal gevulde flessen en hij vindt voor dit gewicht een normale verdeling met gemiddelde 585,2 gram en spreiding 12,8 gram. Vervolgens weegt hij een groot aantal lege flessen die hij van zijn klanten teruggekregen heeft, met de bijbehorende sluitcapsules. Voor dit gewicht vindt hij eveneens een normale verdeling met gemiddelde 228,3 gram en spreiding 11,3 gram.

Bereken, in de veronderstelling, dat het gewicht van de olie in de flessen normaal verdeeld is en onafhankelijk is van het gewicht van fles + sluiting:

- de kansverdeling van het gewicht aan olie in de fles,
- het percentage der flessen olie, die minder dan 350 gram olie bevat-

ten,

- c) de covariantie van het gewicht van een gevulde fles en dat van de olie die erin zit.
- d) Indien men nu, door zorgvuldiger vullen van de flessen, de spreiding van fles + inhoud, bij gelijkblijvend gemiddelde, zoveel kleiner wil maken, dat slechts 2,5% der flessen minder dan 350 gram olie bevat, tot welk bedrag zal men dan de spreiding der gevulde flessen moeten verminderen?

(Opgave 67 uit het Opgavenboekje; zie voor de oplossing blz. 106-107 daarvan.)

OPMERKING. Deze opgave geeft een typerend beeld van enigszins primitieve toepassingen van de statistiek, zoals deze in de praktijk vaak voorkomen. Het aantal metingen wordt zo groot verondersteld, dat het verschil tussen de uit deze metingen gevonden gemiddelden en de daardoor geschatte verwachtingen (eigenlijk: gemiddelden over "alle" flessen olie van die fabrikant) te verwaarlozen is. Hetzelfde geldt voor de spreidingen. De normaliteit van de kansverdelingen is een onderstelling, die onderzocht wordt door de vorm van de frequentieverdeling van de metingen na te gaan. Vertoont deze een "redelijke" gelijkenis met een normale verdelingsdichtheid, dan wordt de normaliteit aanvaard. De uitkomsten van dit soort berekeningen dienen steeds als slechts bij benadering juist te worden beschouwd. De nauwkeurigheid van dergelijke benaderingen zullen wij later leren berekenen.

STELLING 11.6 (Centrale limietstelling voor onderling onafhankelijke stochastische grootheden).

Laat $x_1, x_2, \dots$ een rij stochastisch onafhankelijke grootheden voorstellen. Laat μ_i en σ_i verwachting en spreiding van x_i ($i = 1, 2, \dots$) zijn en laat verder

$$(11.7) \quad \delta_i^3 \stackrel{\text{def}}{=} E|x_i - \mu_i|^3 \quad *)$$

en

*)

Dit is het zgn. absolute derde moment.

$$(11.8) \quad \left\{ \begin{array}{l} \underline{x}(n) \stackrel{\text{def}}{=} \sum_{i=1}^n \underline{x}_i, \\ \mu(n) \stackrel{\text{def}}{=} \sum_{i=1}^n \mu_i, \\ \sigma^2(n) \stackrel{\text{def}}{=} \sum_{i=1}^n \sigma_i^2, \\ \delta^3(n) \stackrel{\text{def}}{=} \sum_{i=1}^n \delta_i^3. \end{array} \right.$$

Indien nu de voorwaarde

$$(11.9) \quad \lim_{n \rightarrow \infty} \frac{\delta(n)}{\sigma(n)} = 0$$

vervuld is, dan is de grootheid

$$(11.10) \quad \underline{x}(n)^* \stackrel{\text{def}}{=} \frac{\underline{x}(n) - \mu(n)}{\sigma(n)}$$

voor $n \rightarrow \infty$ asymptotisch $N(0,1)$ -verdeeld,*) d.w.z. dat voor iedere a geldt:

$$(11.11) \quad \lim_{n \rightarrow \infty} P(\underline{x}(n)^* \leq a) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^a e^{-\frac{1}{2}u^2} du.$$

OPMERKINGEN.

1. De grootheid $\underline{x}(n)^*$ is de gestandaardiseerde van $\underline{x}(n)$, want deze heeft $\mu(n)$ als verwachting en $\sigma(n)$ als spreiding.
2. Het bewijs van deze stelling valt buiten het kader van deze cursus. Het kan bijv. gevonden worden in H. CRAMER, *Mathematical methods of statistics*, Princeton Univ. Press, Princeton 1946, p.215 e.v.

Behalve in de hier gegeven vorm bestaan er nog verschillende andere vormen van de centrale limietstelling; de verschillen hebben dan betrekking op de voorwaarden, waaronder de asymptotische normaliteit geldt.

Als $\underline{x}_1, \underline{x}_2, \dots$ alle dezelfde verdeling bezitten, geldt de volgende stelling:

*) In plaats van deze exacte wijze van uitdrukken zullen wij ook vaak de volgende gebruiken: " $\underline{x}(n)$ is asymptotisch normaal verdeeld" (waarbij dus nog geen standaardisering toegepast is).

STELLING 11.7 (Centrale limietstelling voor onderling onafhankelijke identiek-verdeelde stochastische grootheden).

Zijn $x_1, x_2, \dots$ onderling onafhankelijk verdeeld volgens dezelfde verdeling met eindige verwachting μ en variantie σ^2 , dan is

$$(11.12) \quad \frac{\sum_{i=1}^n x_i - n\mu}{\sigma\sqrt{n}} = \frac{\frac{1}{n} \sum_{i=1}^n x_i - \mu}{\sigma/\sqrt{n}},$$

voor $n \rightarrow \infty$ asymptotisch $N(0,1)$ -verdeeld.

OPMERKINGEN.

1. Het bewijs van stelling 11.7 is bijv. te vinden in H. CRAMÉR, *Mathematical methods of statistics*, p.214-215.
2. Voor het speciale geval dat $x_1, x_2, \dots$ alle dezelfde verdeling bezitten, kan men (11.9) vervangen door

$$(11.9') \quad \delta^3 \stackrel{\text{def}}{=} E|x_i - \mu|^3 \text{ is eindig.}$$

Bewijs dit. Stelling 11.7 volgt dus, als aan (11.9') voldaan is, uit stelling 11.6. Deze voorwaarde kan echter vervangen worden door de zwakkere voorwaarde van eindigheid van μ en σ^2 .

Een belangrijke toepassing van 11.7 is de volgende.

STELLING 11.8 (Stelling van De Moivre en Laplace).

Bezit $\underline{x}$ een binomiale verdeling met parameters n en p , d.w.z. is voor $x = 0, 1, \dots, n$

$$(11.13) \quad P(\underline{x} = x) = \binom{n}{x} p^x q^{n-x} \quad (q = 1 - p),$$

dan is $\underline{x}$ voor $n \rightarrow \infty$ asymptotisch normaal verdeeld. Dus voor iedere a geldt:

$$(11.14) \quad \lim_{n \rightarrow \infty} P\left(\frac{\underline{x} - np}{\sqrt{npq}} \leq a\right) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^a e^{-\frac{1}{2}u^2} du.$$

BEWIJS. Op blz. 77 hebben wij er reeds gebruik van gemaakt, dat $\underline{x}$ beschouwd kan worden als de som van n onafhankelijk verdeelde stochastische grootheden x_i , die "het aantal successen bij het i -de experiment" voorstellen. Met behulp daarvan worden daar de formules $\mu(\underline{x}) = np$ en $\sigma^2(\underline{x}) = npq$

afgeleid. Daarmede is de stelling echter teruggebracht tot een speciaal geval van stelling 11.7, met $\mu = \mu(\underline{x}_i) = p$ en $\sigma^2 = \sigma^2(\underline{x}_i) = pq$. $\square$

OPMERKINGEN.

1. Stelling 11.8 volgt ook uit stelling 11.6. Voor de grootheden $\underline{x}_i$ geldt n.l. dat δ^3 (zie (11.9')) eindig is. (Bewijs dit.)
2. De stelling werd door DE MOIVRE en LAPLACE langs geheel andere weg afgeleid. DE MOIVRE gaf de stelling niet voor de verdelingsfunctie, maar bewees de analoge stelling voor ieder der termen apart. Deze krijgt dan de vorm

$$(11.15) \quad P(\underline{x}=\underline{x}) \approx \frac{1}{\sqrt{2\pi npq}} e^{-\frac{1}{2} \frac{(x-np)^2}{npq}},$$

waarin $\approx$ betekent: "asymptotisch, voor $n \rightarrow \infty$, gelijk aan".

De oorspronkelijke publicaties zijn:

A. DE MOIVRE, *Approximatio ad summam terminorum binomii* $(a+b)^n$ in *seriem expansi*, 1733.

P.S. DE LAPLACE, *Théorie analytique des probabilités*, 1812.

Op grond van deze stelling kan de *normale verdeling* gebruikt worden als *benadering van de binomiale*. Zijn de parameters van de binomiale verdeling n en p , dan is volgens (9.23) en (9.24) de verwachting np en de spreiding $\sqrt{npq}$. De normale verdeling met dezelfde verwachting en spreiding wordt nu de *aangepaste normale verdeling* genoemd.

In fig. 11.1 is, voor $p=0,1$ en $p=0,36$ en voor $n=4,9,25,100$ zowel de binomiale verdeling als de aangepaste normale verdeling geschetst. De x -schaal is daarbij in ieder der 8 figuren zo gekozen, dat de normale verdeling dezelfde vorm bezit (de spreidingen zijn dus (in cm) gelijk gemaakt; hoe wordt dit bereikt?).

Stelling 11.8 wordt geïllustreerd door het feit, dat de benadering duidelijk beter wordt bij toenemende n . Tevens blijkt uit de figuren, dat bij dezelfde waarde van n de benadering voor $p=0,36$ beter is dan voor $p=0,1$. De aanpassing is op zijn best voor $p = \frac{1}{2}$, hetgeen men reeds kan vermoeden op grond van het feit, dat de binomiale verdeling dan symmetrisch is, evenals de normale.

De laatste twee stellingen zijn van groot belang voor het berekenen van bepaalde kansen. Wij geven hiervan twee voorbeelden.

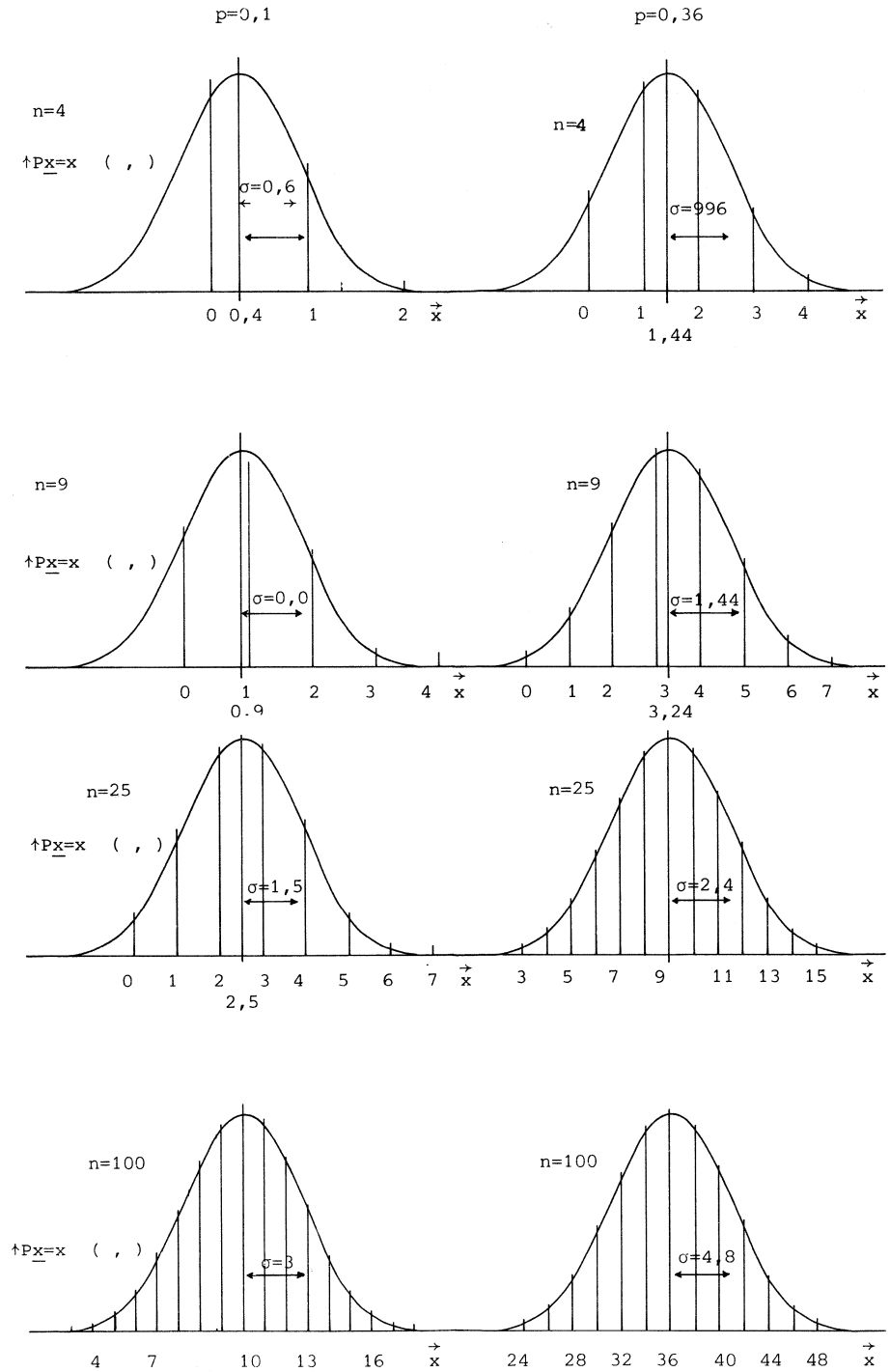


Fig.11.1. Convergentie van de binomiale verdeling naar de normale.

VOORBEELD 11.1. Op blz. 12 hebben wij gezien, dat in 1954 in 6 wijken van Amsterdam tezamen op een totaal van 1927 geboorten 974 jongens voorkwamen. Wij vragen ons nu af, hoe groot de kans op een zo groot of nog groter aantal jongens zou zijn als bij ieder der geboorten de kans op een jongen gelijk is aan $\frac{1}{2}$.

Het aantal jongens, $\underline{x}$, heeft dan een binomiale verdeling met $n = 1927$ en $p = \frac{1}{2}$, dus met

$$E\underline{x} = np = 963,5 \quad \text{en} \quad \sigma(\underline{x}) = \sqrt{npq} = \frac{1}{2}\sqrt{1927} = 21,9.$$

De gevraagde kans is

$$P(\underline{x} \geq 974; n = 1927, p = \frac{1}{2}),$$

waarin achter de puntkomma gegevens staan aangegeven. Wij berekenen deze kans nu - bij benadering - met behulp van (11.14) op de volgende wijze (zonder de gegevens over n en p telkens weer te vermelden).

$$P(\underline{x} \geq 974) = P\left(\frac{\underline{x} - np}{\sqrt{npq}} \geq \frac{974 - np}{\sqrt{npq}}\right) \approx P\left(\underline{u} \geq \frac{974 - 963,5}{21,9}\right) = P(\underline{u} \geq 0,48),$$

waarin $\underline{u}$ een $N(0,1)$ -verdeelde grootheid voorstelt en $\approx$ betekent, dat wij met een benadering te maken hebben.

Volgens tabel 1 van de $N(0,1)$ -verdeling is de kans in het laatste lid gelijk aan 0,32. Onze uitkomst is dus

$$P(\underline{x} \geq 974; n = 1927, p = \frac{1}{2}) \approx 0,32.$$

Het belang van kansen van dit type, die overschrijdingskansen genoemd worden, zullen wij later leren kennen.

VOORBEELD 11.2. Van een verpakkingsmachine, die thee in pakjes van nominaal a gram verpakt is uit vroegere waarnemingen bekend dat de werkelijk verpakte hoeveelheid $\underline{x}$ een normale verdeling bezit met spreiding σ gram. De verwachting μ van de verdeling kan men regelen met een instellings-mechanisme.

Volgens de warenwet moet elk pakje minstens het nominale gewicht bevatten. Om te bereiken dat slechts een fractie α (met $\alpha < \frac{1}{2}$) van de pakjes

minder dan a gram bevat zal men μ groter dan a moeten nemen.

Gevraagd wordt met hoeveel procent het overwicht $\mu - a$ verminderd kan worden indien de controle slechts zou eisen dat het gemiddelde van een aselechte steekproef van n pakjes minstens a gram moet bedragen, terwijl weer een kans α toegelaten wordt dat hieraan niet voldaan is?

OPLOSSING. Bij verwachting μ geldt:

$$P(\underline{x} \leq a) = P\left(\frac{\underline{x} - \mu}{\sigma} \leq \frac{a - \mu}{\sigma}\right) = P\left(\underline{u} \leq \frac{a - \mu}{\sigma}\right),$$

waarin $\underline{u}$ een $N(0,1)$ -verdeelde grootheid voorstelt. Deze kans moet gelijk zijn aan α . Als nu ξ_α gedefinieerd wordt door

$$P(\underline{u} \geq \xi_\alpha) = \alpha,$$

dan moet dus gelden

$$\frac{a - \mu}{\sigma} = -\xi_\alpha \quad \text{of} \quad \mu = a + \xi_\alpha \sigma.$$

Het overwicht moet dus $\xi_\alpha \sigma$ bedragen om ervoor te zorgen dat slechts een fractie α van de pakjes minder dan a gram bevat.

Stelt $\bar{x}$ het gemiddelde gewicht van n aselekt gekozen pakjes voor dan is (vergelijk (10.5) en (10.6)) $E\bar{x} = E\underline{x} = \mu$ en

$$\sigma(\bar{x}) = \frac{\sigma(\underline{x})}{\sqrt{n}} = \frac{\sigma}{\sqrt{n}}.$$

Derhalve is

$$P(\bar{x} \leq a) = P\left(\underline{u} \leq \frac{a - \mu}{\sigma} \sqrt{n}\right)$$

en deze kans is α als

$$\frac{a - \mu}{\sigma} \sqrt{n} = -\xi_\alpha \quad \text{of} \quad \mu = a + \xi_\alpha \frac{\sigma}{\sqrt{n}}.$$

Het overwicht moet nu dus $\xi_\alpha \sigma / \sqrt{n}$ bedragen. De besparing in overwicht is dus, in procenten van het oorspronkelijk overwicht

$$100 \frac{\xi_\alpha \sigma - \xi_\alpha \sigma / \sqrt{n}}{\xi_\alpha \sigma} = 100 \left(1 - \frac{1}{\sqrt{n}}\right).$$

Deze uitkomst is dus onafhankelijk van a, σ en α .

OPGAVE 11.5. Een automatische sorteermachine sorteert eieren op gewicht. Het gewicht van de eieren is normaal verdeeld met verwachting $\mu = 42,0$ gr. en spreiding $\sigma = 5,2$ gr. De sorteermachine scheidt de eieren in 5 klassen.

klasse	gewicht in gr.
A	≥ 55
B	$45 \leq B < 55$
C	$35 \leq C < 45$
D	$25 \leq D < 35$
E	< 25

Gevraagd: hoeveel procent van de eieren bevat ieder der klassen?

(Opgave 25 van het Opgavenboekje; oplossing op blz. 70-71 daarvan.)

OPGAVE 11.6. Van een partij assen is de diameter normaal verdeeld met verwachting 14,82 mm en spreiding 0,03 mm. Van een partij boringen is de diameter eveneens normaal verdeeld en wel met verwachting 14,89 mm en spreiding 0,04 mm.

- Bereken welk percentage assen te groot zal zijn voor de boringen indien as en boring aselekt aan elkaar worden toegevoegd.
- Bereken ook op welke waarde de verwachting van de diameter van de boringen (bij gelijkblijvende spreiding en zonder verandering bij de assen) ingesteld moet worden, opdat slechts 1% van de assen te groot voor de boringen zal zijn, indien wederom as en boring aselekt aan elkaar worden toegevoegd.

(Opgave 78 van het Opgavenboekje; oplossing op blz. 117-118 daarvan.)

OPGAVE 11.7. $\underline{x}$ en $\underline{y}$ zijn onderling onafhankelijk $N(0,1)$ -verdeeld. Bereken de kansverdeling van $|\underline{x} - \underline{y}|$.

(Opgave 93 van het Opgavenboekje; oplossing op blz. 137-138 daarvan.)

OPGAVE 11.8. Bewijs de volgende stelling.

Bezit $\underline{x}$ een negatief benomiale verdeling met parameters k en p , d.w.z. is voor $x = k, k+1, \dots$

$$(11.16) \quad P(\underline{x} = x) = \binom{x-1}{k-1} p^k q^{x-k} \quad (q = 1-p),$$

dan is $\underline{x}$ voor $k \rightarrow \infty$ asymptotisch normaal verdeeld.

Voor iedere a geldt:

$$(11.17) \quad \lim_{k \rightarrow \infty} P\left(\frac{p\underline{x} - k}{\sqrt{kq}} \leq a\right) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^a e^{-\frac{1}{2}u^2} du.$$

AANWIJZING. Vergelijk blz. 77-78 en de centrale limietstelling 11.7.

OPGAVE 11.9. Hoe groot is bij benadering de kans, dat in een reeks van onafhankelijke experimenten met ieder een kans 0,36 op succes, het 100-ste succes niet voor 320-ste experiment verkregen wordt.

OPLOSSING.

- a) Met behulp van (11.17), met $p = 0,36$ en $k = 100$; $\underline{x}$ stelt het aantal experimenten tot en met het 100-ste succes voor.

$$P(\underline{x} \geq 320) = P\left(\frac{0,36\underline{x} - 100}{\sqrt{100 \cdot 0,64}} \geq \frac{0,36 \cdot 320 - 100}{\sqrt{100 \cdot 0,64}}\right) \approx P(\underline{u} \geq 1,9) = 0,029.$$

($\underline{u}$ stelt een $N(0,1)$ -verdeelde grootheid voor.)

- b) Met behulp van de binomiale verdeling. Laat, om verwarring te voorkomen, $\underline{y}$ het aantal successen onder de eerste 319 experimenten voorstellen. Als $y < 100$ is, is het aantal experimenten tot en met het 100-ste succes ≥ 320 en omgekeerd. Dus berekenen wij (nu met behulp van stelling 11.8):

$$\begin{aligned} P(\underline{y} \leq 99) &= P\left(\frac{\underline{y} - np}{\sqrt{npq}} \leq \frac{99 - np}{\sqrt{npq}}\right) \approx P\left(\underline{u} \leq \frac{99 - 319 \cdot 0,36}{\sqrt{319 \cdot 0,36 \cdot 0,64}}\right) = \\ &= P(\underline{u} \leq -1,85) = P(\underline{u} \geq 1,85) = 0,032. \end{aligned}$$

OPMERKING. De beide benaderingen geven niet precies dezelfde uitkomst, hetgeen voor 2 verschillende benaderingsmethoden ook niet nodig is. Zou men de beide berekeningen exact uitvoeren, dan zou de uitkomst wel dezelfde zijn.

HOOFDSTUK XII

SCHATTINGSTHEORIE; MEEST AANNEMELIJKE SCHATTINGEN

De schattingstheorie is een belangrijk onderdeel van de statistiek, die wij hier wegens de grote uitgebreidheid van het onderwerp slechts oppervlakkig kunnen behandelen. Wij geven eerst een aantal voorbeelden met een intuïtief voor de hand liggende oplossing en behandelen daarna het principe van een algemene theorie, waarop deze oplossingen gebaseerd kunnen worden.

VOORBEELD 12.1. Onderzoek van de eigenschappen van een serie-product.

Eén der belangrijkste eigenschappen van gloeilampen is hun levensduur. Dit geldt niet alleen voor gloeilampen, doch voor zeer vele andere producten, in het bijzonder voor onderdelen van machines en apparaten. De hier te beschrijven statistische bewerkingen komen verder op ieder gebied voor. Wij nemen hier, om de gedachten te bepalen, gloeilampen als voorbeeld. De levensduur van een gloeilamp kan men bepalen door hem te laten branden tot hij doorbrandt. De levensduur is dan bekend, maar de lamp is stuk. Het is daarom duidelijk, dat men in dat geval gedwongen is de eigenschappen van een bepaald soort gloeilampen steekproefsgewijs te onderzoeken, een volledig onderzoek schiet zijn doel voorbij.

De levensduur van verschillende gloeilampen van één soort is verschillend; dit blijkt bij beproeving. Een geschikt beschrijvingsmodel van deze eigenschap is daarom de levensduur als een stochastische grootheid x te beschouwen, die voor iedere lamp een bepaalde waarde aanneemt. De bij een steekproef van lampen gevonden waarden vormen dan een reeks van stochastisch onafhankelijke waarnemingen van x . Veelal wordt bij dergelijke problemen verondersteld, dat x een (bij benadering) *normale verdeling* bezit. Soms geldt dit niet voor x zelf, maar wel voor $\log x$. Een dergelijke

veronderstelling kan, op grond van de waarnemingen, ook statistisch worden onderzocht. Wij zullen deze onderstelling hier hanteren alsof zij reeds onderzocht en bevestigd is.

Is nu $\underline{x}$ normaal verdeeld, dan zijn twee parameters, n.l. $\mu = E\underline{x}$ en $\sigma^2 = \sigma^2(\underline{x})$, voldoende om de verdeling volledig te karakteriseren en het schattingsprobleem is dan deze twee parameters uit een reeks waarnemingen te schatten.

Laat de steekproef van waarnemingen voorgesteld worden door de onafhankelijke grootheden

$$(12.1) \quad \underline{x}_1, \underline{x}_2, \dots, \underline{x}_n,$$

alle met dezelfde verdeling als $\underline{x}$, dan hebben wij (vgl. (10.5)) reeds gezien, dat voor het *gemiddelde*

$$(12.2) \quad \bar{\underline{x}} = \frac{1}{n} \sum_{i=1}^n \underline{x}_i$$

geldt:

$$(12.3) \quad E\bar{\underline{x}} = \mu.$$

Het ligt daarom voor de hand om $\bar{\underline{x}}$ het gemiddelde der waarnemingen, als schatting voor μ te gebruiken.

Is een verdelingsfunctie $F(x)$ behalve van x ook nog afhankelijk van een of meer andere niet stochastische variabelen $\theta_1, \theta_2, \dots$, kort aangegeven met θ en $F(x; \theta)$, dan noemt men deze andere variabelen $\theta_1, \theta_2, \dots$ *parameters*. Als een verdelingsfunctie afhankelijk is van een parameter θ dan natuurlijk ook de bijbehorende kansverdeling, notatie $P(\underline{x}=x; \theta)$, of $P_\theta(\underline{x}=x)$, of verdelingsdichtheid, notatie $f(x; \theta)$. Twee verschillende parameterwaarden karakteriseren dus i.h.a. twee verschillende verdelingsfuncties, kansverdelingen of verdelingsdichtheden.

De verzameling Θ van alle mogelijke waarden van deze parameters wordt de *parameter ruimte* genoemd.

Als ϕ een functie is op Θ dan worden de functiewaarden $\phi(\theta)$ ook *parameters* genoemd.

De normale verdelingen worden gekarakteriseerd door de parameters μ en σ , de parameter ruimte bestaat uit alle paren μ en σ met $-\infty < \mu < \infty$ en $\sigma > 0$; ook σ^2 en $\frac{\sigma}{|\mu|}$, bijvoorbeeld, zijn parameters.

De binomiale verdelingen worden gekarakteriseerd door de parameters

n en p , de parameterruimte bestaat uit alle paren n en p met n een natuurlijk getal en $0 < p < 1$ terwijl bijvoorbeeld np en $np(1-p)$ ook parameters worden genoemd.

De hypergeometrische verdelingen hebben parameters M , N en n waarbij M , N en n natuurlijke getallen zijn met $n \leq M+N$.

Een functie $\underline{t} = t(\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n)$ van waarnemingen $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ van een stochastische grootheid $\underline{x}$ met verdelingsfunctie $F(x; \theta)$ die waarden t aanneemt in de parameterruimte van θ wordt een *schatter* van de onbekende parameter θ genoemd.

Een gerealiseerde waarde t van een schatter $\underline{t}$ wordt een *schatting* van de onbekende parameter θ genoemd.

DEFINITIE 12.1. Een schatter $\underline{t}$ van een onbekende parameter θ wordt *zuiver* ^{*)} genoemd als zijn verwachting gelijk is aan deze onbekende parameter

$$(12.4) \quad E \underline{t} = E_{\theta} \underline{t} = \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} t(x_1, \dots, x_n) f(x_1, \dots, x_n; \theta) dx_1 dx_2 \dots dx_n = \theta.$$

DEFINITIE 12.2. Een schatter van een onbekende parameter θ wordt *asymptotisch raak* ^{**)} genoemd, als hij voor een onbegrensd stijgend aantal waarnemingen stochastisch convergeert naar die onbekende parameter:

$$(12.5) \quad \lim_{n \rightarrow \infty} P(|\underline{t} - \theta| > \varepsilon) = 0 \quad \text{voor iedere } \varepsilon > 0.$$

Volgens stelling 10.2 is $\bar{\underline{x}}$ zowel zuiver als asymptotisch raak met betrekking tot μ , ook als $\bar{\underline{x}}$ niet normaal verdeeld is.

Vervolgens gaan wij ons beraden over een *schatter voor de variantie* σ^2 . Daartoe overwegen wij, dat bij de berekening van een schatting in zekere zin de steekproef van waarnemingen in de plaats wordt gesteld van de gehele onbekende kansverdeling. De waarnemingen van de steekproef zijn verder alle gelijkwaardig. Zijn de aangenomen waarden: $x_1, x_2, \dots, x_n$, dan kunnen wij - op grond van deze heuristische redentatie - een nieuwe stochastische grootheid $\underline{x}$ beschouwen, die deze n waarden met gelijke kansen aanneemt:

*) Engels: unbiased.

**) " : consistent.

$$(12.6) \quad P(\underline{x}' = x_i) = \frac{1}{n} \quad (i = 1, 2, \dots, n).$$

Voor deze stochastische grootheid geldt

$$(12.7) \quad E\underline{x}' = \frac{1}{n} \sum_{i=1}^n x_i = \bar{x},$$

de bij deze steekproef behorende waarde van $\bar{x}$, de schatter van μ . Dus: de verwachting van $\underline{x}'$ kan dienen als schatter voor de verwachting van $\underline{x}$. Naar analogie hiervan kunnen wij verwachten, dat de variantie van $\underline{x}'$ ons een schatter voor de variantie van $\underline{x}$ zal geven.

Volgens de definitie van een variantie is:

$$\sigma^2(\underline{x}') = E(\underline{x}' - E\underline{x}')^2 = E(\underline{x}' - \bar{x})^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2,$$

zodat wij de grootheid

$$(12.8) \quad \underline{s}'^2 \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2,$$

beschouwd als schatter voor σ^2 , op zijn merites kunnen onderzoeken.

Voorlopig verstaan wij daaronder, dat wij nagaan of deze schatter zuiver en asymptotisch raak is. Eerst de *zuiverheid*. De verwachting van $\underline{s}'^2$ berekenen wij als volgt.

$$\begin{aligned} \sum_{i=1}^n (x_i - \bar{x})^2 &= \sum_{i=1}^n (x_i - \mu + \mu - \bar{x})^2 = \\ &= \sum_{i=1}^n (x_i - \mu)^2 - 2(\bar{x} - \mu) \sum_{i=1}^n (x_i - \mu) + n(\bar{x} - \mu)^2. \end{aligned}$$

Hierin is

$$\sum_{i=1}^n (x_i - \mu) = n(\bar{x} - \mu),$$

zodat

$$\sum_{i=1}^n (x_i - \bar{x})^2 = \sum_{i=1}^n (x_i - \mu)^2 - n(\bar{x} - \mu)^2.$$

Daar

$$E(x_i - \mu)^2 = \sigma^2 \quad \text{en} \quad E(\bar{x} - \mu)^2 = \sigma^2(\bar{x}) = \frac{\sigma^2}{n},$$

geldt

$$(12.9) \quad E_{\underline{s}}'^2 = \frac{n-1}{n} \sigma^2.$$

De schatting $\underline{s}'^2$ is dus geen zuivere schatter van σ^2 , maar

$$(12.10) \quad \underline{s}^2 \stackrel{\text{def}}{=} \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2$$

is dat wel. Voor praktische toepassingen wordt daarom veelal $\underline{s}^2$ als schatter van σ^2 gebruikt.

NUMERIEK VOORBEELD 12.1. Van 10 lampen van een bepaald type, aselect gekozen uit de productie, is de levensduur (in uren) gemeten. De uitkomsten zijn vermeld in de eerste kolom van tabel 12.1.*)

Tabel 12.1. Berekening van gemiddelde en variantie van een reeks van waarnemingen.

x_i	$x_i - 1000$	$(x_i - 1000)^2$
982	- 18	324
996	- 4	16
1047	47	2209
980	- 20	400
1057	57	3249
1025	25	625
990	- 10	100
1058	58	3364
994	- 6	36
1036	36	1296
totaal	165	11619

Om de berekening te vereenvoudigen kiezen wij eerst een z.g. "voorlopig gemiddelde", d.i. een getal x_0 (het doet er niet veel toe welk precies),

*)

Deze gegevens zijn niet aan de praktijk ontleend.

dat vermoedelijk in de buurt van het werkelijke gemiddelde $\bar{x}$ ligt. In dit geval ligt het voor de hand $x_0 = 1000$ te nemen. In de tweede kolom staat $x_i - 1000$ berekend, in de derde $(x_i - 1000)^2$. Nu is

$$\sum_{i=1}^n x_i = \sum_{i=1}^n (x_i - x_0) + nx_0,$$

dus

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n (x_i - x_0) + x_0.$$

In ons geval

$$\bar{x} = \frac{1}{10} \cdot 165 + 1000 = 1016,5.$$

Verder is (vgl. blz.102)

$$\sum_{i=1}^n (x_i - \bar{x})^2 = \sum_{i=1}^n (x_i - x_0)^2 - n(x_0 - \bar{x})^2.$$

OPMERKING 12.1. Het argument, dat $\underline{s}^2$ een zuivere schatter van σ^2 is en $\underline{s}'^2$ niet, verliest zijn waarde, indien men niet σ^2 maar σ wenst te schatten, hetgeen in de praktijk gewoonlijk het geval is. Immers beschouwt men $\underline{s}$ als schatter van σ , dan is $\underline{s}$ weer onzuiver, zoals uit (9.21) volgt. Vullen wij daarin $\underline{s}$ in voor $\underline{x}$, dan komt er

$$E\underline{s}^2 > (E\underline{s})^2,$$

dus

$$\sigma^2 > (E\underline{s})^2 \quad \text{of} \quad \sigma > E\underline{s},$$

zodat de verwachting van $\underline{s}$ kleiner is dan de te schatten parameter σ .

Het gebruik van $\underline{s}^2$ in plaats van $\underline{s}'^2$ berust dan ook in feite op traditie. Het bovenstaande geldt nog steeds zonder de veronderstelling van normaliteit. Voor normaal verdeelde $\underline{x}$ kan de factor

$$c_n \stackrel{\text{def}}{=} \frac{\sigma}{E\underline{s}}$$

berekend worden. De schatter $c_n \underline{s}$ is dan een zuivere schatter voor σ .

Zowel s als s' zijn *asymptotisch raak*; voor $n \rightarrow \infty$ zijn zij equivalent, daar dan $\frac{n-1}{n} \rightarrow 1$. Beschouw nu

$$s'^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 - (\bar{x} - \mu)^2.$$

Volgens stelling 10.2 convergeert $\bar{x}$ stochastisch naar μ , dus $|\bar{x} - \mu|$ naar 0, dus ook $(\bar{x} - \mu)^2$ naar 0. Deze term kan dus verwaarloosd worden: behoudens een willekeurig kleine kans wordt hij zelf willekeurig klein. Noemen wij verder

$$y_i \stackrel{\text{def}}{=} (x_i - \mu)^2, \quad (i = 1, 2, \dots, n),$$

dan is

$$E y_i = \sigma^2, \quad (i = 1, 2, \dots, n)$$

en, wederom volgens stelling 10.2 convergeert nu

$$\frac{1}{n} \sum_{i=1}^n y_i = \frac{1}{n} \sum_{i=1}^n (x_i - \mu)^2$$

stochastisch naar de verwachting, dus naar σ^2 *). Daarmede is echter het gestelde (zij het niet volledig exact) bewezen. In ons geval

$$\sum_{i=1}^n (x_i - \bar{x})^2 = 11619 - 10(16,5)^2 = 8896,5.$$

Derhalve is

$$s^2 = \frac{1}{9} \cdot 8896,5 = 988,5, \quad s = 31,4.$$

Willen wij nu bijv. schatten, hoe groot het percentage lampen is, dat een levensduur < 950 uur bezit, dan berekenen wij

$$\frac{1016,5 - 950}{31,4} = 2,12$$

en zoeken dit getal op in de tabel van de $N(0,1)$ -verdeling. De uitkomst is 0,017, dus 1,7%. Van de steekproef is dan ook geen enkele uitkomst zo laag.

*) Volgens stelling 10.2 moet de voorwaarde vervuld zijn dat $\sigma^2(y_i)$ eindig is; deze voorwaarde kan echter vermeden worden.

VOORBEELD 12.2. Kurken tellen

Een aardig historisch voorbeeld van een eenvoudig schattingsprobleem is het volgende ^{*)}. In een kurkenfabriek, die miljoenen kurken in voorraad had, waarvan dagelijks duizenden in kleinere en grotere partijen moesten worden afgeleverd, waren voortdurend een aantal arbeiders bezig met de kand kurken te tellen. Dit was een vermoeiende en zenuwslopende bezigheid, waarmee het bovendien nog niet mogelijk was aan het einde van het jaar de voorraad op te nemen. Men wenste daarom een kurkentel-apparaat te ontwikkelen.

Aangezien het bij aflevering van partijen kurken niet op een paar meer of minder aankomt, kon Ir. VAN ETTINGER de oplossing in statistische richting zoeken. Zijn methode komt hierop neer, dat uit een zak kurken van één maat een steekproef genomen wordt, waarvan het totale gewicht ($\underline{g}$) en het aantal (n) bepaald wordt. Het gemiddelde gewicht van één kurk is dan $\underline{g}/n$. Is nu het gewicht van de gehele zak kurken gelijk aan G , dan is

$$\frac{nG}{\underline{g}}$$

een schatter voor het aantal kurken uit de zak.

Deze methode werd bovendien als volgt gemechaniseerd. Een weegschaal werd vervaardigd, die in evenwicht is als het gewicht aan kurken op de éne schaal 99 maal zo groot is als dat op de andere. Men zet nu de zak kurken op de ene schaal en legt daaruit zoveel kurken op de andere, dat de weegschaal in evenwicht is. Het totale aantal kurken is dan, bij benadering, gelijk aan 100 maal het aantal op de tweede schaal overgebrachte kurken. De tijdsbesparing is enorm, de investering gering en de arbeiders zijn trots op de eenvoudige wijze waarop zij kurken tellen. Het systeem is, in die fabriek, reeds bijna 30 jaar in gebruik.

VOORBEELD 12.3. Enquête

Wij beschouwen vervolgens een enquête, die erop gericht is de kennis van een bepaalde groep mensen te onderzoeken. Voor de eenvoud vatten wij één vraag in het oog, die juist of onjuist beantwoord kan worden. Deze vraag wordt voorgelegd aan een steekproef van n personen, aselekt genomen uit de te onderzoeken groep. In statistisch jargon wordt deze te onderzoeken

^{*)}

J. VAN ETTINGER, *Statistiek en productiviteit*, Statistica 6 (1952), pp. 19-28, pag.20.

groep personen de *populatie* genoemd. Als model nemen wij aan, dat een fractie p van de personen uit deze groep de vraag goed zou beantwoorden en de overige (fractie $q = 1 - p$) fout. Het gaat er nu om deze fracties te schatten op grond van de antwoorden van de n personen van de steekproef.

Kiezen wij aselekt één persoon uit de populatie, dan is de kans, dat wij een juist antwoord verkrijgen gelijk aan p . Bestaat de populatie uit N personen, dan bevinden zich daaronder dus Np , die een juist antwoord zouden geven en Nq , die de vraag onjuist zouden beantwoorden. Hieruit zijn er nu n aselekt en zonder teruglegging genomen en wij bevinden ons dus in de situatie van voorbeeld 7 van hoofdstuk V. Het aantal juiste antwoorden in de steekproef van omvang n , dat wij met $\underline{x}$ aangeven, heeft dus een hypergeometrische verdeling.

$$(12.11) \quad P(\underline{x}=\underline{x}) = \frac{\binom{Np}{\underline{x}} \binom{Nq}{n-\underline{x}}}{\binom{N}{n}},$$

met, volgens (9.26),

$$(12.12) \quad E\underline{x} = np,$$

dus

$$(12.13) \quad E \frac{\underline{x}}{n} = p.$$

De fractie juiste antwoorden in de steekproef is dus een zuivere schatter voor de onbekende fractie p . De vraag van asymptotisch gedrag doet zich hier niet op natuurlijke wijze voor, daar $n > N$ niet op kan treden. Voor $n = N$ vindt men p reeds exact.

Wat wel kan gebeuren is, dat N niet bekend is. Volgens (12.12) kan $\underline{x}/n$ dan toch als schatter gebruikt worden. Is $N \gg n$, dan is volgens (5.10) de kansverdeling van $\underline{x}$ bij benadering *binomiaal*:

$$(12.14) \quad P(\underline{x}=\underline{x}) \approx \binom{n}{\underline{x}} p^{\underline{x}} q^{n-\underline{x}},$$

waarbij dan (12.12) geldig blijft. Aan de schatter verandert dus niets. Volgens stelling 10.3 (de theoretische wet van de grote aantallen) is de schatter in het geval van een binomiale verdeling *asymptotisch raak*. (In het hier beschouwde geval impliceert $n \rightarrow \infty$ dat ook $N \rightarrow \infty$; is p bijv. de onbekende kans op succes bij een experiment, dat willekeurig vaak herhaald kan

worden, dan doet deze complicatie zich niet voor.)

VOORBEELD 12.4. Schieten op de kermis

Een enthousiast kermisganger schiet in de schiettent op Goudse pijpen. Bij ieder schot kan hij alleen zien of hij het pijpje, waar hij op mikte, geraakt heeft of niet, maar als hij mist kan hij niet zien, waar zijn kogel terecht gekomen is. Hij kan bij een volgend schot dus geen gebruik maken van het resultaat van het vorige schot om zijn richten te verbeteren. Als model kunnen wij daarom stellen, dat hij bij ieder schot dezelfde kans p op succes heeft en dat de schoten onafhankelijk zijn.

De schutter neemt zich nu voor net zo lang door te gaan met schieten tot hij k pijpjes stukgeschoten heeft. Hiervoor heeft hij $\underline{x}$ schoten nodig. Gevraagd wordt uit deze gegevens een schatter voor p af te leiden.

Op het eerste gezicht ligt het in dit geval wellicht voor de hand om, net als bij het vorige voorbeeld, de fractie successen, dus hier $k/\underline{x}$, als schatter voor p te nemen.

De verdeling van $\underline{x}$ is in dit geval *negatief binomiaal* (formule (7.7)) en volgens (9.25) geldt

$$E\underline{x} = \frac{k}{p} \quad \text{dus} \quad p = \frac{k}{E\underline{x}},$$

hetgeen een verdere rechtvaardiging in schijnt te houden.

Dit is echter niet geheel juist, daar uit deze formules niet volgt dat de schatter zuiver is. Daartoe zou n.l. moeten gelden

$$E \frac{k}{\underline{x}} = kE \frac{1}{\underline{x}} = p,$$

dus zou

$$E \frac{1}{\underline{x}} = \frac{1}{E\underline{x}}$$

moeten zijn en dit is niet het geval. In de oplossing van opgave 85 van het Opgavenboekje (blz.129-130) wordt bewezen, dat wel geldt

$$(12.15) \quad E \frac{k-1}{\underline{x}-1} = p,$$

zodat $k-1/\underline{x}-1$ wel een zuivere schatter van p is. Deze correctie van -1 in teller en noemer wordt intuïtief min of meer begrijpelijk als men overweegt,

dat het laatste schot bij deze proefopzet altijd een succes is, daar het k -de succes het sein tot beëindiging van de serie is. Laat men nu dit laatste schot geheel weg, dan gaat de serie over in een reeks van $\underline{x} - 1$ schoten met $k - 1$ successen, die nu overal in de serie geplaatst kunnen zijn, juist als bij een binomiale verdeling het geval is.

Voor $k \rightarrow \infty$ zijn de beide schatters $k/\underline{x}$ en $k-1/\underline{x}-1$ equivalent. Nu kan (vgl. blz. 77 en 78) $\underline{x}$ beschouwd worden als de som van k onafhankelijke grootheden $\underline{x}_i$, die alle dezelfde (negatief binomiale) verdeling (met $k = 1$) bezitten. Volgens stelling 10.2 convergeert dan het gemiddelde

$$\bar{\underline{x}} = \frac{1}{k} \sum_{i=1}^k \underline{x}_i$$

van deze grootheden stochastisch naar hun verwachting, die $\frac{1}{p}$ is. Met andere woorden $\underline{x}/k$ zal, voor voldoende grote k , willekeurig dicht tot $\frac{1}{p}$ naderen. Dit betekent echter dat $k/\underline{x}$ willekeurig dicht bij p komt. Beide schatters zijn dus asymptotisch raak.

OPMERKING. Voor $k = 1$ is het niet mogelijk de beschreven correctiemethode toe te passen, daar de teller van de schatter dan $= 0$ zou worden. Voor dit geval is geen praktisch bruikbare (men kan gemakkelijk nagaan dat de stochastische grootheid gedefinieerd door $\underline{t} = 1$ als $\underline{x} = 1$ en $\underline{t} = 0$ als $\underline{x} > 1$ een zuivere schatter voor p is) zuivere schatter voor p beschikbaar.

De verwachting van $1/\underline{x}$ voor $k = 1$ laat zich als volgt berekenen.

$$E \frac{1}{\underline{x}} = \sum_{x=1}^{\infty} \frac{1}{x} p q^{x-1} = \frac{p}{q} \sum_{x=1}^{\infty} \frac{q^x}{x} = - \frac{p}{q} \ln(1-q).$$

Voor $p = q = \frac{1}{2}$ wordt dit $\ln(2) = 0,69$. De verwachting van de schatter verschilt dan dus aanzienlijk van de te schatten waarde.

Wij behandelen vervolgens de elementen van de *theorie der meest aannemelijke schatters*, waarvan de grondslag gelegd is door R.A. FISHER. Hoewel deze theorie niet de enige schattingstheorie is, is het momenteel toch wel een van de belangrijkste en wij zullen ons tot deze theorie beperken.

Het principe van de theorie zetten wij uiteen voor continue verdelingen. Voor discrete is het analoog, met kansen in plaats van verdelingsdichtheden.

Laat

$$(12.16) \quad x_1, x_2, \dots, x_n$$

n (al of niet onafhankelijke) waarnemingen voorstellen, die een simultane verdeling bezitten, met verdelingsdichtheid

$$(12.17) \quad f(x_1, x_2, \dots, x_n; \theta_1, \theta_2, \dots, \theta_h),$$

waarin $\theta_1, \theta_2, \dots, \theta_h$ onbekende parameters voorstellen, die uit de waarnemingen geschat moeten worden.

Stellen nu $x_1, x_2, \dots, x_n$ de waarden voor, die in werkelijkheid gevonden zijn, zodat het dus getallen zijn, geen variabelen meer, dan is (12.17) een functie van de onbekende parameters θ_j ($j = 1, 2, \dots, h$).

Deze functie wordt de *aannemelijkheidsfunctie* of kortweg *aannemelijkheid* genoemd en aangegeven als

$$(12.18) \quad L(\theta_1, \theta_2, \dots, \theta_h; x_1, x_2, \dots, x_n) \stackrel{\text{def}}{=} f(x_1, x_2, \dots, x_n; \theta_1, \theta_2, \dots, \theta_h).$$

Gewoonlijk (doch niet voor alle verdelingsdichtheden) bezit L een maximum; het punt, waarin dit wordt aangenomen, wordt aangegeven met $\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_h$ en deze waarden - die functies van de x_i ($i = 1, 2, \dots, n$) zijn - worden de *meest aannemelijke schattingen* van $\theta_1, \theta_2, \dots, \theta_h$ genoemd. De $\hat{\theta}_1, \hat{\theta}_2, \dots, \hat{\theta}_h$ zijn, opgevat als stochastische grootheden, de *meest aannemelijke schatters* van $\theta_1, \theta_2, \dots, \theta_h$.

Toelichting aan voorbeelden

VOORBEELD 12.1. Normale verdeling (zie ook blz. 104).

In dit voorbeeld heeft x een normale verdeling met onbekende μ en σ . De waarnemingen zijn $x_1, x_2, \dots, x_n$. Daar deze stochastisch onafhankelijk zijn, is

$$(12.19) \quad L(\mu, \sigma^2; x_1, x_2, \dots, x_n) = \left(\frac{1}{\sigma\sqrt{2\pi}} \right)^n \prod_{i=1}^n e^{-\frac{(x_i - \mu)^2}{2\sigma^2}} = \\ = \left(\frac{1}{\sigma\sqrt{2\pi}} \right)^n e^{-\frac{1}{2} \frac{\sum_{i=1}^n (x_i - \mu)^2}{\sigma^2}}.$$

In deze vorm beschouwen wij σ^2 (en niet σ) als onbekende parameter, daar dit de berekeningen vereenvoudigt. Om het maximum van L te vinden, bij

gegeven waarden der x_i , differentiëren wij, ter verdere vereenvoudiging, niet L zelf, maar de logaritme van L naar μ en naar σ^2 .

$$(12.20) \quad \log L = -\frac{n}{2} \log 2\pi - \frac{n}{2} \log \sigma^2 - \frac{1}{2} \frac{\sum_{i=1}^n (x_i - \mu)^2}{\sigma^2},$$

$$(12.21) \quad \frac{\partial \log L}{\partial \mu} = \frac{\sum_{i=1}^n (x_i - \mu)}{\sigma^2},$$

$$(12.22) \quad \frac{\partial \log L}{\partial \sigma^2} = -\frac{n}{2\sigma^2} + \frac{1}{2} \frac{\sum_{i=1}^n (x_i - \mu)^2}{\sigma^4}.$$

Stellen wij deze afgeleiden gelijk aan 0, dan verkrijgen wij dus 2 vergelijkingen voor $\hat{\mu}$ en $\hat{\sigma}^2$, en wel

$$(12.23) \quad \sum_{i=1}^n (x_i - \hat{\mu}) = 0$$

en

$$(12.24) \quad \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \hat{\mu})^2.$$

Uit de eerste volgt

$$(12.25) \quad \hat{\mu} = \frac{1}{n} \sum_{i=1}^n x_i = \bar{x} \quad (\text{vgl. (12.2)})$$

en vervolgens uit de tweede

$$(12.26) \quad \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2 = s'^2 \quad (\text{vgl. (12.8)}).$$

De meest aannemelijke schatters zijn dus $\bar{x}$ en s'^2 en komen dus in dit geval overeen met de boven op intuïtieve wijze gekozenen.

OPMERKINGEN.

12.2. Indien wij niet σ^2 , maar σ zelf als onbekende parameter gekozen hadden, zou dezelfde uitkomst verkregen zijn. Dit is een speciaal geval van een algemene stelling, inhoudende dat de meest aannemelijke schatter van een continue 1,1-functie van een onbekende parameter verkregen wordt door

dezelfde functie te nemen van de meest aannemelijke schatter van die parameter. Dus: is $\hat{\theta}$ de meest aannemelijke schatter van θ , dan is $\phi(\hat{\theta})$ de meest aannemelijke schatter van $\phi(\theta)$ als ϕ een continue en ondubbelzinnig omkeerbare functie is.

12.3. Wij zien tevens, dat meest aannemelijke schatters niet zuiver behoeven te zijn. Dit volgt trouwens reeds uit opmerking 12.2, daar zuiverheid, zoals wij in opmerking 12.1 gezien hebben, bij transformatie van de parameter verloren kan gaan.

VOORBEELD 12.3. Binomiale verdeling (zie ook blz. 112).

Is x binomiaal verdeeld met onbekende p , maar bekende n , dan geldt us

$$(12.27) \quad L(p; x) = P(x=x; p) = \binom{n}{x} p^x q^{n-x}.$$

Dus

$$(12.28) \quad \log L = \log \binom{n}{x} + x \log p + (n-x) \log (1-p).$$

Differentiëren naar p geeft

$$(12.29) \quad \frac{d \log L}{dp} = \frac{x}{p} - \frac{n-x}{1-p},$$

zodat wij, om $\hat{p}$ te vinden, de vergelijking

$$(12.30) \quad \frac{x}{\hat{p}} = \frac{n-x}{1-\hat{p}}$$

moeten oplossen. Dit geeft

$$(12.31) \quad \hat{p} = \frac{x}{n}.$$

De meest aannemelijke schatter voor p is dus $\frac{x}{n}$.

OPMERKING 12.4.

Ook hier vinden wij dus de vroeger reeds vermelde schatter. Voor de hypergeometrische verdeling is dit niet het geval; daarbij treden complicaties op, o.a. doordat Np in dat geval geheel moet zijn, zodat p niet alle waarden tussen 0 en 1 kan bezitten. Wij zullen dit geval hier buiten beschouwing laten.

VOORBEELD 12.4. Negatief binomiale verdeling (zie ook blz. 113).

Is $\underline{x}$ negatief binomiaal verdeeld met onbekende p , maar bekende k , dan is

$$(12.32) \quad L(p; \underline{x}) = \binom{x-1}{k-1} p^k q^{x-k},$$

dus

$$(12.33) \quad \log L = \log \binom{x-1}{k-1} + k \log p + (x-k) \log(1-p)$$

en

$$(12.34) \quad \frac{d \log L}{dp} = \frac{k}{p} - \frac{x-k}{1-p}.$$

De op te lossen vergelijking is nu dus

$$(12.35) \quad \frac{k}{p} = \frac{x-k}{1-p},$$

waaruit volgt

$$(12.36) \quad \hat{p} = \frac{k}{x},$$

zodat de meest aannemelijke schatter wordt $\frac{k}{x}$.

VOORBEELD 12.5. Hoogwaterstanden; exponentiële verdeling

Indien men de frequentieverdeling van hoogwaterstanden waargenomen in Hoek van Holland, beschouwt, blijkt dat deze, voor zoverre het niet te kleine waarden betreft, bij goede benadering exponentieel verdeeld zijn. Bedoeld wordt, dat voor iedere niet te lage waarde x_0 geldt

$$(12.37) \quad P(\underline{x} \geq x | \underline{x} \geq x_0) = e^{-\lambda(x-x_0)} \quad (x > x_0)^*.$$

Geldt dit voor één waarde x_0 , dan ook voor iedere grotere waarde $x'_0 > x_0$. Immers uit (12.37) volgt

$$e^{-\lambda(x-x_0)} = \frac{P(\underline{x} \geq x \wedge \underline{x} \geq x_0)}{P(\underline{x} \geq x_0)}.$$

*)

P.J. WEMELSFELDER, *Wetmatigheden in het optreden van stormvloed*, "De Ingenieur", 3 maart 1939.

Daar $x \geq x_0$ is de teller gelijk aan $P(\underline{x} \geq x)$. De noemer geven wij aan met P_0 . Dan geldt dus

$$(12.38) \quad P(\underline{x} \geq x) = P_0 e^{-\lambda(x-x_0)}.$$

Voor een willekeurige waarde $x'_0 > x_0$ geldt dus

$$(12.39) \quad P(\underline{x} \geq x'_0) = P_0 e^{-\lambda(x'_0-x_0)}.$$

Voor $x \geq x'_0$ is nu

$$P(\underline{x} \geq x | \underline{x} \geq x'_0) = \frac{P(\underline{x} \geq x \wedge \underline{x} \geq x'_0)}{P(\underline{x} \geq x'_0)}.$$

Hierin is de teller weer gelijk aan $P(\underline{x} \geq x)$ en voor de noemer geldt (12.39). Dus is

$$(12.40) \quad P(\underline{x} \geq x | \underline{x} \geq x'_0) = \frac{P_0 \cdot e^{-\lambda(x-x_0)}}{P_0 \cdot e^{-\lambda(x'_0-x_0)}} = e^{-\lambda(x-x'_0)} \quad (x \geq x'_0).$$

Dit betekent, dat de keuze van het "beginpunt" x_0 (of x'_0) de vorm van de verdeling - en in het bijzonder de waarde van λ - niet beïnvloedt, hetgeen het mogelijk maakt de lage waarden van x , waar de verdeling niet aan (12.37) voldoet, uit te schakelen zonder dat de betrekkelijk willekeurige keuze van het beginpunt veel invloed op de uitkomst uitoefent.

Deze eigenschap van de exponentiële verdeling, dat de vorm van de voorwaardelijke verdeling onder de voorwaarde $\underline{x} \geq x_0$ voor iedere x_0 dezelfde is, is daarom voor het statistische onderzoek van de hoogwaterstanden aan de Nederlandse kust - en daarmee voor de bepaling van dijkhoogten, die een voldoende bescherming tegen overstromingen bieden - van veel belang. Een verklaring voor het experimenteel vastgestelde feit, dat de verdeling der hoogwaterstanden exponentieel is, is (nog?) niet bekend.

Dat de verdeling deze vorm bij goede benadering bezit kan experimenteel vastgesteld worden door gebruik te maken van logaritmisch grafiekenpapier. Noemen wij het linkerlid van (12.37) P , dan geldt

$$(12.41) \quad \log P = -\lambda(x-x_0).$$

Zetten wij dus $\log P$ verticaal uit tegen $x - x_0$ horizontaal, dan wordt een rechte lijn verkregen, zoals in fig. 12.1 geschetst is.

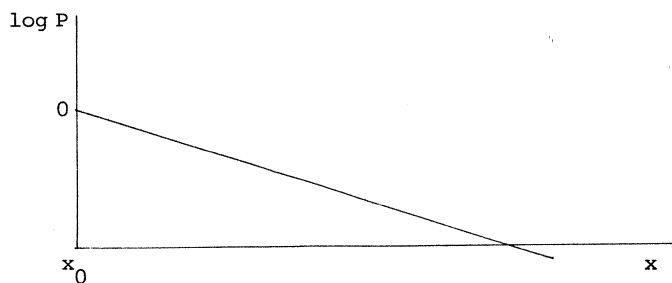


Fig. 12.1. Verdeling van hoogwaterstanden
in Hoek van Holland.

Deze lijn is, uiteraard, onbekend en moet uit waarnemingen geschat worden. Daartoe wordt hieronder de meest aannemelijke schatter voor λ bepaald. Beschouwen wij daartoe alle hoogwaterstanden boven x_0 , die in een bepaalde periode zijn opgetreden, dan kunnen deze (indien bepaalde voorzorgen genomen worden, die hier onbesproken blijven) beschouwd worden als onafhankelijke waarnemingen van een stochastische grootheid $\underline{x}$, die de door (12.37) gegeven verdeling bezit. De verdelingsdichtheid, steeds onder de voorwaarde $\underline{x} \geq x_0$, is

$$(12.42) \quad f(x) = F'(x) = \frac{d(1-P)}{dx} = \lambda e^{-\lambda(x-x_0)},$$

dus als $x_1, x_2, \dots, x_n$ de waarnemingen zijn, is

$$(12.43) \quad L(\lambda; x_1, x_2, \dots, x_n) = \prod_{i=1}^n \lambda e^{-\lambda(x_i - x_0)} = \lambda^n e^{-\lambda \sum_{i=1}^n (x_i - x_0)}.$$

Verder

$$(12.44) \quad \log L = n \log \lambda - \lambda \sum_{i=1}^n (x_i - x_0)$$

en

$$(12.45) \quad \frac{d \log L}{d\lambda} = \frac{n}{\lambda} - \sum_{i=1}^n (x_i - x_0).$$

Derhalve is

$$(12.46) \quad \hat{\lambda} = \frac{n}{\sum_{i=1}^n (x_i - x_0)} = \frac{1}{\bar{x} - x_0},$$

dus $\hat{\lambda}^{-1}$ is het gemiddelde der waarnemingen boven x_0 verminderd met x_0 .

Nu leidt (12.46) niet tot een zuivere schatter. Beschouwen wij n.l. de schatter (12.46) weer als een stochastische grootheid, d.w.z.

$$\hat{\lambda} = \frac{1}{\bar{x} - x_0},$$

terwijl steeds de voorwaarde $\underline{x} \geq x_0$, dus ook $\underline{x}_i \geq x_0$ ($i = 1, 2, \dots, n$) geldt, dan wordt in §2 van de appendix aangetoond dat

$$(12.47) \quad E\left(\frac{1}{\bar{x} - x_0} \mid \underline{x}_1 \geq x_0, \underline{x}_2 \geq x_0, \dots, \underline{x}_n \geq x_0\right) = \frac{1}{n-1} \lambda,$$

zodat

$$\frac{n-1}{n} \frac{1}{\bar{x} - x_0} = \frac{n-1}{\sum_{i=1}^n (\underline{x}_i - x_0)}$$

wel een zuivere schatter van λ is.

Voor $n \rightarrow \infty$ zijn de twee schatters equivalent. Verder is

$$E(\hat{\lambda}^{-1} \mid \underline{x}_1 \geq x_0, \dots, \underline{x}_n \geq x_0) = E(\bar{x} - x_0 \mid \underline{x}_1 \geq x_0, \dots, \underline{x}_n \geq x_0) = E(\underline{x} - x_0 \mid \underline{x} \geq x_0)$$

en de (voorwaardelijke) verdelingsdichtheid van $\underline{x}$ is (zie (12.42))

$$f(x) = \lambda e^{-\lambda(x-x_0)},$$

zodat

$$E(\underline{x} - x_0 \mid \underline{x} \geq x_0) = \lambda \int_{x_0}^{\infty} (x - x_0) \cdot e^{-\lambda(x-x_0)} dx = \lambda \int_0^{\infty} v e^{-\lambda v} dv.$$

Nu is echter (partieel integreren)

$$\lambda \int_0^{\infty} v e^{-\lambda v} dv = \left[-v e^{-\lambda v} \right]_0^{\infty} + \int_0^{\infty} e^{-\lambda v} dv = \left[-\frac{1}{\lambda} e^{-\lambda v} \right]_0^{\infty} = \frac{1}{\lambda}.$$

Dus is $1/\hat{\lambda}$ wel een zuivere schatter van $1/\lambda$. De schatter $1/(\bar{x} - x_0)$ (en dus ook $(n-1)/(n(\bar{x} - x_0))$) is wel asymptotisch raak. Uit stelling 10.2 volgt n.l. dat $\bar{x} - x_0$ stochastisch convergeert naar $E(\underline{x} - x_0 \mid \underline{x} \geq x_0) = 1/\lambda$ en dus convergeert $\hat{\lambda} = 1/(\bar{x} - x_0)$ stochastisch naar λ .

OPMERKING 12.5. Bij het werkelijke onderzoek van deze waterstanden werden niet alle hoogwaterstanden gebruikt, doch slechts een in overleg met het K.N.M.I. op grond van meteorologische overwegingen uitgekozen groep, behorende bij bepaalde depressies van een potentieel gevaarlijk karakter.

VOORBEELD 12.6. Regressie; theorie der kleinste kwadraten

Wij beschouwen, als laatste voorbeeld, een stochastische grootheid $\underline{y}$, waarvan de verdeling afhangt van een andere, niet stochastische grootheid x . Zo kan $\underline{y}$ bijv. de in een week door een koe geleverde hoeveelheid melk voorstellen (die van week tot week en van koe tot koe verschilt) en x de leeftijd (in jaren) van de koe. Of $\underline{y}$ kan de lengte van de remweg van een auto voorstellen en x de snelheid bij het begin van het remmen. In deze beide gevallen kan men x zelf kiezen - in het eerste geval door een koe van de gewenste leeftijd uit te zoeken, in het tweede door bij de gewenste snelheid te gaan remmen -, terwijl $\underline{y}$ waargenomen wordt.

In beide gevallen gaat het erom te onderzoeken hoe de kansverdeling van $\underline{y}$ van x afhangt.

Wij beschouwen hier alleen het eenvoudigste geval, n.l. dat alleen de verwachting van $\underline{y}$ van x afhankelijk is en wel lineair:

$$(12.48) \quad E(\underline{y};x) = \alpha + \beta x,$$

waarin α en β onbekende parameters zijn, die geschat moeten worden. Verder veronderstellen wij, dat de verdeling van $\underline{y}$, bij gegeven x , normaal verdeeld is met een onbekende spreiding σ , die niet van x afhangt.

Op dit probleem gaan wij nu de methode der meest aannemelijke schatters toepassen. De waarnemingen, die verondersteld worden stochastisch onafhankelijk te zijn, zijn nu van de vorm

$$(12.49) \quad (x_1, \underline{y}_1), (x_2, \underline{y}_2), \dots, (x_n, \underline{y}_n),$$

daar iedere waar te nemen waarde $\underline{y}_i$ bij een bepaalde waarde x_i van x behoort. De waarnemingsuitkomsten $(x_i, \underline{y}_i)$ ($i = 1, 2, \dots, n$) kunnen als punten in een vlak worden uitgezet, zoals in fig. 12.2 gedaan is.

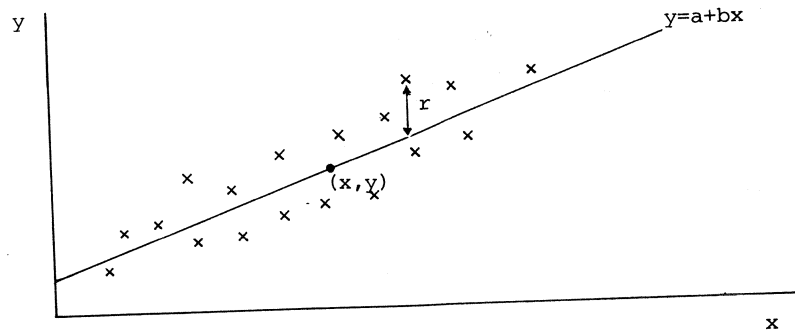


Fig. 12.2. Lineaire regressie van y op x.

De verdelingsdichtheid van y_i is

$$(12.50) \quad f(y_i; x_i) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2} \frac{\{y_i - E(y_i; x_i)\}^2}{\sigma^2}}.$$

Wegens de onafhankelijkheid der y_i is

$$(12.51) \quad L(\alpha, \beta, \sigma^2; y_1, \dots, y_n) = \prod_{i=1}^n f(y_i; x_i)$$

en dus (vgl. (12.48))

$$(12.52) \quad \log L = -\frac{n}{2} \log 2\pi - \frac{n}{2} \log \sigma^2 - \frac{1}{2\sigma^2} \sum_{i=1}^n (y_i - \alpha - \beta x_i)^2.$$

Voor de schatting van α en β is alleen de laatste term van belang. Deze geven wij (onder weglaten van de factor $-1/2\sigma^2$) aan met Q .

Daar $\log L$ gemaximaliseerd moet worden, dient Q geminimaliseerd te worden. Voor de eenvoud van de berekeningen voeren wij het gemiddelde $\bar{x}$ der x_i in en schrijven α' def $\alpha + \beta\bar{x}$, zodat

$$(12.53) \quad a + \beta x_i = \alpha' + \beta(x_i - \bar{x}).$$

Nu wordt Q :

$$(12.54) \quad Q = \sum_{i=1}^n \{y_i - \alpha' - \beta(x_i - \bar{x})\}^2.$$

Door uitvermenigvuldigen van de kwadraten kan hiervoor geschreven worden

$$Q = \sum_{i=1}^n (y_i - \alpha')^2 - 2\beta \sum_{i=1}^n (y_i - \alpha') (x_i - \bar{x}) + \beta^2 \sum_{i=1}^n (x_i - \bar{x})^2.$$

Voor de tweede term kan geschreven worden

$$- 2\beta \sum_{i=1}^n y_i (x_i - \bar{x}) + 2\beta\alpha' \sum_{i=1}^n (x_i - \bar{x})$$

en daar $\sum_{i=1}^n (x_i - \bar{x}) = 0$ valt de laatste term weg, zodat verkregen wordt

$$(12.55) \quad Q = \sum_{i=1}^n (y_i - \alpha')^2 + \beta^2 \sum_{i=1}^n (x_i - \bar{x})^2 - 2\beta \sum_{i=1}^n y_i (x_i - \bar{x}).$$

Q bestaat dus uit 2 stukken, waarvan het éne alleen α' en het andere alleen β bevat, zodat deze afzonderlijk geminimaliseerd kunnen worden. Beschouw eerst de eerste term, die α' bevat.

$$\begin{aligned} \sum_{i=1}^n (y_i - \alpha')^2 &= \sum_{i=1}^n (y_i - \bar{y} + \bar{y} - \alpha')^2 = \\ &= \sum_{i=1}^n (y_i - \bar{y})^2 + 2(\bar{y} - \alpha') \sum_{i=1}^n (y_i - \bar{y}) + n(\bar{y} - \alpha')^2 = \\ &= \sum_{i=1}^n (y_i - \bar{y})^2 + n(\bar{y} - \alpha')^2, \end{aligned}$$

daar

$$\sum_{i=1}^n (y_i - \bar{y}) = 0.$$

De eerste der overgebleven termen bevat α' niet en de tweede term wordt minimaal als voor α' de waarde $\bar{y}$ ingevuld wordt. Derhalve is de meest aanmerkelijke schatting van α' :

$$(12.56) \quad \alpha' \stackrel{\text{def}}{=} \bar{\alpha}' = \bar{y}.$$

De laatste twee termen van Q in (12.55), die β bevatten, worden, zoals bekend, tezamen minimaal als voor β de waarde

$$\sum_{i=1}^n \frac{y_i (x_i - \bar{x})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

ingevuld wordt, dus is

$$(12.57) \quad b \stackrel{\text{def}}{=} \hat{\beta} = \frac{\sum_{i=1}^n y_i (x_i - \bar{x})}{\sum_{i=1}^n (x_i - \bar{x})^2}.$$

De schattingen voor α' en β zijn dus verkregen en uit (12.48) tezamen met (12.53) volgt nu als schatting y voor $E(y; x)$

$$(12.58) \quad y = a' + b(x - \bar{x}) = a + bx \quad (a \stackrel{\text{def}}{=} a' - b\bar{x}),$$

i.h.b.

$$y_i = a' + b(x_i - \bar{x}) = a + bx_i, \quad i = 1, 2, \dots, n.$$

y wordt de bij x behorende *regressiewaarde* genoemd; (12.58) stelt een rechte lijn voor, die door het punt $(\bar{x}, \bar{y})$ gaat (het "zwaartepunt" van de "puntenwolk" in fig. 12.2) en die de (geschatte) *regressielijn* van y op x genoemd wordt. De "werkelijke" regressielijn wordt door (12.48) gegeven.

Uit (12.52) valt nu ook de meest aannemelijke schatting van σ^2 te berekenen, na in deze vorm voor α en β de reeds gevonden schattingen a en b ingevuld te hebben. Immers voor iedere waarde van σ^2 wordt $\log L$ maximaal voor minimale waarde van

$$(12.59) \quad Q_0 = \sum_{i=1}^n (y_i - a - bx_i)^2.$$

Zo gaat $\log L$ dus over in

$$- \frac{n}{2} \log 2\pi - \frac{n}{2} \log \sigma^2 - \frac{Q_0}{2\sigma^2}$$

en de afgeleide hiervan naar σ^2 is

$$- \frac{n}{2\sigma^2} + \frac{Q_0}{2\sigma^4}.$$

Stellen wij deze vorm gelijk aan 0, dan volgt daaruit de schatting $\hat{\sigma}^2$ van σ^2 :

$$(12.60) \quad \hat{\sigma}^2 = \frac{Q_0}{n}.$$

Vervolgens onderzoeken wij de *zuiverheid* der verkregen schatters. Eerst $\underline{a}'$.

$$\begin{aligned}
 E\bar{a}' &= E\bar{y} = \frac{1}{n} \sum_{i=1}^n E(y_i; x_i) = \frac{1}{n} \sum_{i=1}^n \{\alpha' + \beta(x_i - \bar{x})\} = \\
 &= \alpha' + \frac{\beta}{n} \sum_{i=1}^n (x_i - \bar{x}) = \alpha',
 \end{aligned}$$

daar

$$\sum_{i=1}^n (x_i - \bar{x}) = 0.$$

Derhalve is $\bar{a}'$ een zuivere schatter van α' . Vervolgens $\bar{b}$.

$$\begin{aligned}
 E \sum_{i=1}^n y_i (x_i - \bar{x}) &= \sum_{i=1}^n (x_i - \bar{x}) E y_i = \sum_{i=1}^n (x_i - \bar{x}) \{\alpha' + \beta(x_i - \bar{x})\} = \\
 &= \alpha' \sum_{i=1}^n (x_i - \bar{x}) + \beta \sum_{i=1}^n (x_i - \bar{x})^2 = \beta \sum_{i=1}^n (x_i - \bar{x})^2.
 \end{aligned}$$

Dus

$$E\bar{b} = \frac{\beta \sum_{i=1}^n (x_i - \bar{x})^2}{\sum_{i=1}^n (x_i - \bar{x})^2} = \beta,$$

zodat $\bar{b}$ een zuivere schatter van β blijkt te zijn.

Uit deze resultaten volgt, dat ook $\bar{a}$ een zuivere schatter van α is, immers (zie (12.58) en (12.53)):

$$E\bar{a} = E(\bar{a}' - \bar{b}\bar{x}) = \alpha' - \beta\bar{x} = \alpha.$$

De berekening van $E\hat{\sigma}^2$, die iets ingewikkelder is, zullen wij hier niet geven. De uitkomst is

$$(12.61) \quad E\hat{\sigma}^2 = \frac{n-2}{n} \sigma^2,$$

zodat $\hat{\sigma}^2$ geen zuivere schatter is, maar

$$(12.62) \quad \underline{s}^2 \stackrel{\text{def}}{=} \frac{\underline{Q}_0}{n-2} = \frac{\sum_{i=1}^n (y_i - \bar{a} - \bar{b}x_i)^2}{n-2}$$

wel.

De schatters zijn ook *asymptotisch raak*, doch het bewijs daarvoor laten wij achterwege.

OPMERKINGEN.

12.6. De hier toegepaste methode komt, meetkundig, daarop neer dat die lijn $y = a + bx$ in fig. 12.2 gekozen wordt, waarvoor de som van de kwadraten van de afstanden der waargenomen punten tot die lijn, in verticale richting gemeten, minimaal is. Deze afstanden worden *residuen* genoemd. Eén residu is in fig. 12.2 door de letter r aangegeven. Het bij (x_i, y_i) behorende residu is

$$(12.63) \quad r_i \stackrel{\text{def}}{=} y_i - a - bx_i,$$

dus

$$(12.64) \quad Q_0 = \sum_{i=1}^n r_i^2.$$

De methode ontleent daaraan de naam: *methode der kleinste kwadraten*. Onder deze naam is zij reeds zeer lang bekend als de basis van de zgn. *foutentheorie*, die veel toepassing vindt op het gebied van de natuur- en sterrekunde, ook voor veel gecompliceerdere problemen dan de hier behandelde.

12.7. Soms zijn ook de x_i stochastisch. De methode kan dan overanderd toegepast worden, daar men dan gebruik kan maken van de voorwaardelijke verdelingen der y_i onder de voorwaarden $x_i = x_i$ ($i = 1, 2, \dots, n$); deze voorwaardelijke verdelingen moeten dan voldoen aan de boven voor de y_i geformuleerde voorwaarden. Dit is bijv. het geval als x en y een 2-dimensionale normale verdeling bezitten, waaruit de waarnemingspunten (x_i, y_i) een steekproef vormen. Voorbeeld: de lengtes van aselekt gekozen vaders en hun oudste zonen. Naast de regressie van y op x kan dan ook de regressie van x op y onderzocht worden, die een andere uitkomst geeft. Wij gaan hier niet nader op in.

12.8. Ten onrechte wordt veelal ondersteld, dat de methode der kleinste kwadraten alleen voor normaal verdeelde y_i bruikbaar zou zijn. Uit het bovenstaande blijkt, dat bij de bewijzen der zuiverheid geen gebruik is gemaakt van die veronderstelling, zodat zij algemener gelden. Wel geldt speciaal in het normale geval, dat de methode der kleinste kwadraten ook de meest aannemelijke schattingen geeft. In andere gevallen kunnen de volgens de beide methoden verkregen schattingen verschillend zijn.

Nauwkeurigheid van schatters

Een schatter wordt uit waarnemingen berekend. Deze waarnemingen worden als

stochastische grootheden beschouwd, zodat niet alleen zij, maar ook de schatter een kansverdeling bezit. De uitkomst, die verkegen wordt door bepaalde waargenomen waarden in de formule van de schatter te substitueren - een concrete waarde van de schatter - kan zelf als één waarneming van die schatter beschouwd worden.

De kansverdeling van een schatter is een volledige beschrijving van de eigenschappen van die schatter, waarvan bijzondere eigenschappen, zoals zuiverheid en asymptotische raakheid speciale aspecten zijn. Een andere, in het voorafgaande nog niet beschouwde eigenschap is de *nauwkeurigheid* van een schatter, die grofweg door zijn *spreiding* gegeven wordt.

VOORBEELD 12.1 (zie blz. 115). Spreiding van het gemiddelde van een steekproef

Wij beschouwen bij dit voorbeeld een steekproef van n waarnemingen

$$(12.1) \quad \underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$$

uit een normale verdeling met verwachting μ en variantie σ^2 . Als schatter voor μ namen wij

$$(12.2) \quad \bar{\underline{x}} = \frac{1}{n} \sum_{i=1}^n \underline{x}_i,$$

hetgeen ook (vergelijk (12.25)) de meest aannemelijke schatter is. Volgens stelling 11.5 is $\bar{\underline{x}}$ zelf ook normaal verdeeld met (vergelijk (10.5) en (10.6))

$$(12.65) \quad E\bar{\underline{x}} = \mu \quad \text{en} \quad \sigma^2(\bar{\underline{x}}) = \frac{\sigma^2}{n}.$$

Voor de spreiding van $\bar{\underline{x}}$ geldt dus

$$(12.66) \quad \sigma(\bar{\underline{x}}) = \frac{\sigma}{\sqrt{n}}$$

en men kan bewijzen (hetgeen wij hier niet zullen doen), dat er in het geval van een steekproef uit een normale verdeling geen enkele zuivere schatter van μ bestaat, die een kleinere spreiding bezit. Men noemt daarom $\bar{\underline{x}}$ de *meest nauwkeurige zuivere schatter* van μ .

De grootheid

$$(12.67) \quad \bar{\underline{x}}^* = \frac{\bar{\underline{x}} - \mu}{\sigma} \sqrt{n},$$

de gestandaardiseerde van $\bar{x}$, is dus $N(0,1)$ -verdeeld, zodat, volgens tabel 1 bijv. geldt:

$$P(|\bar{x}^*| \leq 1,96) = 0,95$$

of, na omrekenen van de ongelijkheid tussen de haken,

$$(12.68) \quad P\left(\bar{x} - 1,96 \frac{\sigma}{\sqrt{n}} \leq \mu \leq \bar{x} + 1,96 \frac{\sigma}{\sqrt{n}}\right) = 0,95.$$

Bij het numerieke voorbeeld van blz. 108 vonden wij $\bar{x} = 1016,5$. De ongelijkheid tussen de haken van (12.68) gaat, bij substitutie van deze waarde, over in

$$1016,5 - 1,96 \frac{\sigma}{\sqrt{10}} \leq \mu \leq 1016,5 + 1,96 \frac{\sigma}{\sqrt{10}},$$

waarmee wij nog niet veel opschieten als σ onbekend is. Nemen wij even aan, dat σ op grond van vroegere waarnemingen bekend is en bijv. 33 bedraagt, dan vinden wij na invullen van deze waarde

$$(12.69) \quad 996 \leq \mu \leq 1037.$$

Deze uitspraak, waarbij de onbekende parameter μ tussen twee grenzen ingesloten wordt, kan juist of onjuist zijn. Op grond van (12.68) zou men de neiging kunnen hebben te zeggen, dat er een kans 0,95 is, dat hij juist is. Strikt genomen is dit niet waar: (12.69) is òfwel juist (en dan is de kans, dat hij juist is, gelijk aan 1) òfwel onjuist (en dan is die kans gelijk aan 0). Maar wèl volgt uit (12.68), dat de gevolgde *methode* een kans 0,95 bezit op een juiste uitkomst. Beperkt men dus zijn beschouwingen niet tot één geval (zoals de uitkomst (12.69)), maar beschouwt men de methode als zodanig, of een lange reeks toepassingen daarvan, dan kan men zeggen, dat de kans op een juiste uitspraak 0,95 bedraagt.

In algemene vorm wordt (12.68):

$$(12.70) \quad P\left(\bar{x} - u_{\frac{1}{2}\alpha} \frac{\sigma}{\sqrt{n}} \leq \mu \leq \bar{x} + u_{\frac{1}{2}\alpha} \frac{\sigma}{\sqrt{n}}\right) = 1 - \alpha,$$

waarin $u_{\frac{1}{2}\alpha}$ die waarde van u in tabel 1 voorstelt, waarvoor $k_r = \frac{1}{2}\alpha$ is.

Het interval tussen de haken van (12.70) wordt een *betrouwbaarheidsinterval* voor μ genoemd, $1 - \alpha$ heet de *betrouwbaarheidscoefficiënt* daarvan

(dit is dus de kans, dat de methode tot een juiste uitspraak zal leiden) en α heet de *onbetrouwbaarheid*.

Meestal is σ niet bekend. In eerste instantie substitueert men dan veelal een schatter van σ , bijv. $\underline{s} = \sqrt{\underline{s}^2}$ (zie (12.10)) voor σ . Zoals wij op blz. 109 hebben gezien is deze schatter niet zuiver: gemiddeld wordt σ door $\underline{s}$ onderschat. Dit doet vermoeden, dat de grenzen, waartussen μ op deze wijze ingesloten wordt wel eens te dicht bij elkaar kunnen komen, of ook dat de kans op een verkeerde uitspraak bij toepassing van deze methode groter dan α zal zijn. Dit is inderdaad juist en men kan de substitutie van $\underline{s}$ voor σ dan ook slechts als een benadering beschouwen. Voor grote steekproeven ($n \rightarrow \infty$) bestaat dit bezwaar niet. Dan convergeert n.l. $\underline{s}$ stochastisch naar σ . Voor kleine steekproeven bestaat - in het normale geval - een exacte methode, die op de z.g. verdeling van STUDENT berust. Wij gaan daar momenteel niet op in.

Uit (12.70) is (nogmaals) te zien, dat μ in principe met willekeurige nauwkeurigheid bepaald kan worden door n voldoende groot te nemen. Dit is equivalent met het vroeger bewezene, dat $\bar{\underline{x}}$ voor $n \rightarrow \infty$ stochastisch tot μ convergeert.

In de praktijk dient deze stelling echter met voorzichtigheid gehanteerd te worden. Indien men bijv. de lengte van een voorwerp een aantal malen opmeet met een meetinstrument dat niet nauwkeuriger dan in cm afleesbaar is en de werkelijke lengte van het voorwerp is gelijk aan 10,9 cm, dan zullen alle waarnemingen (of vrijwel alle) de waarde 11 cm aangeven en men behoeft er niet op te hopen, dat het gemiddelde tot 10,9 cm zal convergeren. De kansverdeling der waarnemingen is dan ook niet normaal, zelfs niet continu, daar alleen gehele waarden aangenomen kunnen worden. In de praktijk zijn echter waarnemingen altijd van discreet karakter. De moraal daarvan is, dat het in de regel zinloos is het gemiddelde van een aantal metingen in meer decimalen op te geven dan op het meetinstrument afgelezen kunnen worden. Hiertegen wordt vaak gezondigd.

VOORBEELD 12.7. De mediaan in plaats van de verwachting

De *mediaan* van een continu verdeelde stochastische grootte $\underline{x}$ is die waarde $\underline{x}_M$, waarvoor geldt

$$(12.71) \quad P(\underline{x} \leq \underline{x}_M) = P(\underline{x} \geq \underline{x}_M) = \frac{1}{2}.$$

Deze waarde is niet altijd ondubbelzinnig bepaald, in principe kan er een

geheel interval van waarden zijn, die aan (12.71) voldoen. In dat geval wordt dit gehele interval de mediaan genoemd.

Bij discrete verdelingen kan (12.71) op dezelfde wijze gebruikt worden, als er minstens één waarde x_M is, die daaraan voldoet. Dit is niet altijd het geval. De algemene definitie, die voor alle verdelingen van toepassing is, luidt als volgt.

DEFINITIE 12.3. De *mediaan* van een stochastische grootte $\underline{x}$ is de verzameling van alle waarden x_M , die voldoen aan

$$(12.72) \quad \begin{aligned} &P(\underline{x} < x_M) < \frac{1}{2} \text{ en } P(\underline{x} < x_M) < \frac{1}{2} \\ &\text{of} \\ &P(\underline{x} < x_M) = \frac{1}{2} \text{ en } P(\underline{x} < x_M) = \frac{1}{2} \end{aligned}$$

OPMERKING. Deze definitie sluit de mogelijkheid uit dat bij alternatieve diskrete verdelingen, bijvoorbeeld $P(\underline{x}=0) = P(\underline{x}=1) = \frac{1}{2}$, de punten 0 en 1 ook mediaan kunnen zijn.

Voor symmetrische verdelingen (zoals de normale en de binomiale met $p = \frac{1}{2}$) is de mediaan gelijk aan de verwachting. Voor de meeste verdelingen, die scheef naar rechts (resp. naar links) zijn, ligt de mediaan links (resp. rechts) van de verwachting (zie fig. 12.1).

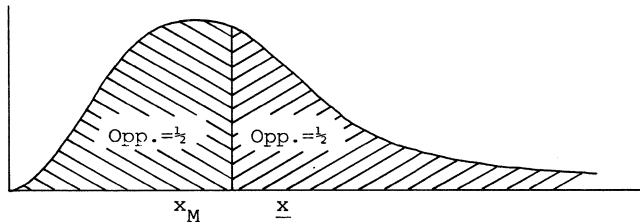


Fig.12.1. Ligging van mediaan en verwachting bij een
continue verdeling, die scheef naar rechts is.

Er is soms reden om de mediaan in plaats van de verwachting te gebruiken als een parameter, die grofweg de ligging van de verdeling aangeeft. Een voorbeeld daarvan treft men bij het onderzoeken van de studieduur bij studenten. Een aantal daarvan studeert nl. in het geheel niet af en heeft dus een oneindige studieduur. Is echter bij een stochastische grootte de kans, dat de waarde ∞ aangenomen wordt, positief, dan is ook de ver-

wachting = ∞ . Voor de mediaan geldt dit niet, die wordt in het geheel niet beïnvloed door het gedrag van de verdeling "in de staarten", maar alleen door het middelste gedeelte. Ook als men alleen die groep studenten beschouwt, die wel hun studie voltooiën, dan behoudt de mediaan nog zekere voordelen boven de verwachting; er zijn nl. altijd studenten, doe om één of andere reden hun studie gedurende lange tijd onderbreken en pas veel later afstuderen en het is vaak moeilijk vast te stellen wat men voor die gevallen de studieduur moet noemen. Dat deze lang is staat vrijwel altijd vast, maar op hoe lang men hem precies moet waarderen is vaak niet bij benadering te zeggen. Deze gevallen hebben echter juist een zeer grote invloed op het gemiddelde (als schatter van de verwachting), maar geen enkele op de mediaan.

Als schatter voor de mediaan x_M wordt gewoonlijk de mediaan van de waarnemingen genomen; dit is de middelste waarneming na rangschikking naar grootte als het aantal oneven is of het gemiddelde der twee middelste waarnemingen als het aantal even is. Notatie: $\underline{x}_m$.

Het onderzoek van de stochastische eigenschappen van deze schatter is vrij moeilijk en valt grotendeels buiten het bestek van de cursus.

Is $\underline{x}$ symmetrisch verdeeld om zijn verwachting (tevens mediaan) μ , dan is $\tilde{x} = \underline{x} - \mu$ symmetrisch om 0, d.w.z. dat $-\tilde{x}$ en $+\tilde{x}$ dezelfde verdeling bezitten. Uit symmetrie-overwegingen volgt dan direct, dat ook de gereduceerde mediaan $\tilde{\underline{x}}_m = \underline{x}_m - \mu$ van een steekproef van waarnemingen symmetrisch om 0 verdeeld is, zodat $\underline{x}_m$ zelf een zuivere schatter van x_M is. Dat de mediaan $\underline{x}_m$ altijd *asymptotisch raak* is als schatter van x_M is vrij gemakkelijk in te zien. Wij bewijzen het alleen voor het geval van een continue verdeling met $f(x) > 0$ voor x in een omgeving van x_M . Laat $\epsilon > 0$ een klein getal zijn en noem

$$p = P(\underline{x} \leq x_M - \epsilon) \quad \text{en} \quad p' = P(\underline{x} \geq x_M + \epsilon),$$

dan is $p < \frac{1}{2}$ en $p' < \frac{1}{2}$. De fractie waarnemingen, die $\leq x_M - \epsilon$ zijn, zal nu, voor $n \rightarrow \infty$, stochastisch convergeren naar p en de fractie, die $\geq x_M + \epsilon$ zijn, naar p' . Behoudens een willekeurig kleine kans zal dus $\underline{x}_m$ tussen $x_M - \epsilon$ en $x_M + \epsilon$ vallen, als n voldoende groot is. Dit geldt echter voor iedere $\epsilon > 0$, waaruit de te bewijzen eigenschap volgt.

In de meeste in de praktijk voorkomende gevallen is de mediaan *onnauwkeuriger* dan het gemiddelde, d.w.z. dat hij een grotere spreiding bezit. In het geval van een normale verdeling geldt, dat $\underline{x}_m$ voor $n \rightarrow \infty$ asymptotisch

normaal verdeeld is met μ als verwachting en

$$(12.73) \quad \sigma(\underline{x}_m) \approx \sigma \sqrt{\frac{\pi}{2n}},$$

als σ de spreiding van de verdeling is. Wij hebben reeds opgemerkt, dat $\bar{x}$ in dat geval de schatter met de kleinste mogelijke spreiding is. Deze is $\sigma/\sqrt{n}$. De limiet van het quotiënt der varianties

$$(12.74) \quad \lim_{n \rightarrow \infty} \frac{\sigma^2(\bar{x})}{\sigma^2(\underline{x}_m)} = \frac{2}{\pi} = 0,63$$

heet de *asymptotische doeltreffendheid* van $\underline{x}_m$.

Ondanks deze kleinere doeltreffendheid wordt bij routinetoepassingen (in het bijzonder controles op lopende-band-productie) van de statistiek in de industrie vaak gebruik gemaakt van de mediaan in plaats van het gemiddelde, ook bij normaal verdeelde grootheden. De reden daarvan is, dat de bepaling van de mediaan geen berekening nodig maakt, zodat deze ook door ongeschoold personeel zonder moeite bepaald kan worden (men gebruikt dan bijv. telkens 5 waarnemingen, waarvan de middelste gemakkelijk te bepalen is). De geringere doeltreffendheid kan dan vaak ondervangen worden door de controle wat vaker te doen plaatsvinden.

VOORBEELD 12.3. (zie blz. 117) Binomiale verdeling

Bij een binomiale verdeling met parameters n en p vonden wij $\underline{x}/n$ als meest aannemelijke schatter. De variantie hiervan volgt direct uit (9.24):

$\sigma^2(\underline{x}) = npq$. Wegens stelling 9.5 is nl.

$$(12.75) \quad \sigma^2\left(\frac{\underline{x}}{n}\right) = \frac{1}{n} \sigma^2(\underline{x}) = \frac{pq}{n}.$$

Ook $\underline{x}/n$ heeft van alle zuivere schatters van p , de kleinste variantie. Deze variantie is weer in de regel onbekend, doch kan uit de waarnemingen geschat worden. Een zuivere schatter ervan is

$$(12.76) \quad \frac{\underline{x}(n-\underline{x})}{n^2(n-1)},$$

zoals uit de volgende berekening blijkt.

Volgens (9.20), (9.23) en (9.24) is

$$E\bar{x}^2 = \sigma^2(\bar{x}) + (E\bar{x})^2 = npq + n^2 p^2.$$

Dus

$$E\bar{x}(n-\bar{x}) = nE\bar{x} - E\bar{x}^2 = n^2 p - npq - n^2 p^2 = n(n-1)pq.$$

Verder weten wij uit stelling 11.8, dat $\bar{x}$ asymptotisch normaal verdeeld is. Dit geldt dan ook voor $\bar{x}/n$. Voor voldoende grote n kan dan ook, analoog aan (12.70) het volgende *benaderde betrouwbaarheidsinterval* voor de onbekende p aangegeven worden:

$$(12.77) \quad P\left(\frac{1}{n} \left\{ \bar{x} - u_{\frac{1}{2}\alpha} \left(\frac{\bar{x}(n-\bar{x})}{n-1} \right)^{\frac{1}{2}} \right\} \leq p \leq \frac{1}{n} \left\{ \bar{x} + u_{\frac{1}{2}\alpha} \left(\frac{\bar{x}(n-\bar{x})}{n-1} \right)^{\frac{1}{2}} \right\}\right) \approx 1 - \alpha.$$

In de literatuur vindt men gewoonlijk onder de wortel de factor n in de noemer in plaats van $n-1$. Daar de formule alleen voor grote waarden van n nauwkeurig is, maakt dit niet veel verschil. De schatter (12.76) is trouwens wel zuiver als schatter van $\sigma^2(\bar{x}/n)$, maar de wortel eruit is weer (vergelijk opmerking 12.1 op blz. 109) een onzuivere schatter van $\sigma(\bar{x}/n)$. In plaats van (12.77) kan een betere benadering opgesteld worden, hetgeen wij echter tot een later hoofdstuk uitstellen, evenals de afleiding van exacte betrouwbaarheidsintervallen voor een onbekende kans.

VOORBEELD 12.6. (zie blz. 122) Lineaire regressie

Als laatste voorbeeld berekenen wij nog de varianties van de *regressie-coëfficiënten* $\underline{a}'$ en $\underline{b}$ uit (12.56) en (12.57). Daar alle $\underline{y}_i$ dezelfde variantie σ^2 bezitten en onafhankelijk verdeeld zijn, volgt uit (12.56) onmiddellijk

$$(12.78) \quad \sigma^2(\underline{a}') = \frac{\sigma^2}{n}.$$

Volgens (12.57) is $\underline{b}$ een lineaire combinatie der $\underline{y}_i$ van de vorm

$$\underline{b} = \sum_{i=1}^n c_i \underline{y}_i,$$

met

$$c_i = \frac{x_i - \bar{x}}{\sum_{i=1}^n (x_i - \bar{x})^2},$$

waarin de x_i bekende (niet stochastische) getallen zijn. Derhalve geldt

$$\sigma^2(\underline{b}) = \sum_{i=1}^n c_i^2 \sigma^2(\underline{y}_i) = \sigma^2 \sum_{i=1}^n c_i^2,$$

dus

$$(12.79) \quad \sigma^2(\underline{b}) = \frac{\sigma^2}{\sum_{i=1}^n (x_i - \bar{x})^2}.$$

Daar verder zowel $\underline{a}'$ als $\underline{b}$ lineaire combinaties van de normaal verdeelde grootheden $\underline{y}_i$ zijn, zijn zij zelf ook normaal verdeeld. Met $\underline{s}^2$ (zie (12.62)) als schatter voor σ^2 kan men op grond daarvan weer benaderde betrouwbaarheidsintervallen voor α' en β opstellen. De afleiding van exacte betrouwbaarheidsintervallen geven wij hier niet. Deze berust weer op de hier nog niet behandelde verdeling van STUDENT.

Bij berekening blijkt de covariantie van $\underline{a}'$ en $\underline{b}$ gelijk aan 0 te zijn. Verder kan men bewijzen, dat $\underline{a}'$ en $\underline{b}$ onafhankelijk verdeeld zijn, hetgeen bij verdere uitwerking - bijv. ter bepaling van een betrouwbaarheidsstrook ook voor de gehele regressielijn - aanzienlijke voordelen geeft. Ook voor deze schattingen geldt, dat zij - althans in het normale geval - de kleinst mogelijke spreiding bezitten.

Tenslotte valt op te merken, dat de *optimale eigenschap* van kleinst mogelijke spreiding, die in het bovenstaande herhaaldelijk is vermeld, niet toevallig ontstaat. Onder vrij algemene voorwaarden bezitten n.l. de meest aannemelijke schatters, althans asymptotisch voor $n \rightarrow \infty$, de eigenschap minimale spreidingen te bezitten en normaal verdeeld te zijn. In vele gevallen is de eerstgenoemde eigenschap ook reeds voor eindige n vervuld, terwijl de laatstgenoemde vaak voorkomt als de waarnemingen normaal verdeeld zijn. Het zijn o.a. deze eigenschappen, die de theorie der meest aannemelijke schatters haar aantrekkelijkheid verlenen.

HOOFDSTUK XIII

TOETSINGSTHEORIE; BINOMIALE TOETSEN

Men kan de meeste elementaire toepassingen van de statistiek globaal in twee groepen verdelen: toepassingen van de schattingstheorie, die in hoofdstuk XII behandeld is, en van de toetsingstheorie, die in dit hoofdstuk ter sprake komt. Het doel van de schattingstheorie is: op grond van steekproeven van waarnemingen tot kwantitatieve schattingen van onbekende parameters te komen; de waarden van deze schattingen staan daarbij dus niet van tevoren vast. De toetsingstheorie heeft ten doel op grond van een steekproef van waarnemingen te onderzoeken of één of meer onbekende parameters bepaalde gegeven waarden bezitten of niet.

In symbolen: is θ een onbekende parameter van een kansverdeling (bijv. de gemiddelde levensduur van een bepaald soort gloeilamp of de overlevingskans van een mug, die op een bepaalde wijze met een insecticide in aanraking wordt gebracht), dan geeft de schattingstheorie een schatter $\underline{t}$ voor θ met een kansverdeling, die min of meer om θ gecentreerd ligt, terwijl de toetsingstheorie erop gericht is te onderzoeken of θ al of niet gelijk is aan een bepaalde gegeven waarde θ_0 .

Er bestaat een nauw verband tussen deze beide theorieën. Zo is het intuïtief reeds duidelijk, dat men, op grond van een gegeven waarnemingsreeks, tot verwerping van de getoetste waarde θ_0 zal dienen over te gaan, indien de voor $\underline{t}$ uit deze waarnemingsreeks gevonden waarde sterk van θ afwijkt. Vindt men bijv. voor een steekproef van 100 lampen een gemiddelde levensduur van 700 uur (met een niet te grote spreiding daar omheen), terwijl men de waarde $\theta_0 = 2000$ wil onderzoeken, dan is deze waarde voor θ niet acceptabel. Dit verband tussen de beide theorieën zal bij de behandeling van de toetsingstheorie duidelijk naar voren komen.

Voor zover het de toepassingen betreft is er echter een duidelijk verschil tussen de twee methoden. Bij vele wetenschappelijke problemen is in

de eerste plaats de toetsing van bepaalde hypothesen van belang en pas in de tweede plaats de schattingstheorie. Indien men bijv. (een klassiek probleem) twee geneesmiddelen op hun merites ter bestrijding van een bepaalde ziekte wil vergelijken, is het in hoofdzaak van belang te onderzoeken of één van de twee beter is dan het andere. Is men er eenmaal van overtuigd, dat middel A beter is dan B, dan zal men, *ceteris paribus*, middel A voortaan prefereren boven B; de grootte van het verschil in geneeskracht is van secundair belang. Wenst men daarentegen bijv. te berekenen welke premie men moet vragen voor een bepaalde verzekering, dan heeft men een kwantitatieve schatting van het risico nodig en het is dan niet genoeg (of zelfs niet van belang) om te weten, dat dit risico groter is dan een bepaald bedrag of groter dan het risico, dat aan een andere situatie verbonden is. De keuze tussen verschillende toepasbare statistische methoden hangt dus, zoals steeds het geval is, af van het *doel* van het onderzoek. Een duidelijke formulering van dit doel leidt gewoonlijk gemakkelijk tot de keuze van een bepaalde methode, terwijl het achterwege laten van deze formulering - een vaak voorkomend euvel, ook bij wetenschappelijke onderzoeken - kan leiden tot een verkeerde keuze.

Wij zullen de toetsingstheorie in algemene bewoordingen beschrijven met daarnaast, ter illustratie, het speciale geval van binomiale toetsen. In het volgende hoofdstuk bespreken wij dan nog enkele vaak gebruikte toetsen.

Terminologie en notatie

Het waarnemingsmateriaal bestaat - naar wij zullen veronderstellen - uit één of meer reeksen van onafhankelijke waarnemingen. Wij geven het aan met

$$(13.1) \quad \underline{w}_1, \underline{w}_2, \dots, \underline{w}_n.$$

Deze stochastische grootheden behoeven niet alle dezelfde verdeling te bezitten; zij kunnen ook meerdimensionaal zijn: ieder ervan kan bijv. een paar van waarnemingen $(\underline{x}_i, \underline{y}_i)$ voorstellen.

VOORBEELD 13.1. Fraaie voorbeelden van het optreden van kansen met bekende waarden vindt men in de erfelijkheidsleer. Het mag wel bekend verondersteld worden, dat aan de proeven van MENDEL een erfelijkheidstheoretisch model ten grondslag ligt, dat bijv. leidt tot de uitkomst, dat bij het kruisen van bepaalde planten één van twee verschijnselen A en B zich voordoet en

dat de kans op A bij iedere proef een bekende waarde (bijv. $\frac{1}{4}$) bezit. De verschijnselen A en B kunnen twee kleuren (van erwten, of bloemen) zijn of iets dergelijks. De n waarnemingen, verkregen bij n onafhankelijke proeven van deze aard, nemen in dit geval de "waarden" A of B aan; men kan dit in cijfers omzetten door bij iedere proef het aantal malen A (dus 1 of 0) als waarneming (w) te beschouwen; (13.1) is dan dus een rij van n énen en nullen.

De toegelaten hypothesen; de parameterruimte

Als de kansverdelingen der w_i volledig bekend waren, zou er geen toetsingsprobleem bestaan. In de regel is er over deze kansverdelingen wel iets bekend, d.w.z. er worden bepaalde veronderstellingen over gemaakt, die als gegeven worden beschouwd. Deze hebben veelal betrekking op stochastische onafhankelijkheid der w_i en vaak ook op de vorm der verdelingen. Zo berusten veel toetsen op de veronderstelling van normaliteit der verdelingen. Wij zullen hier het geval beschouwen, dat de verdelingen der w_i van één of meer onbekende parameters afhangen, maar overigens volledig bekend zijn. Geven wij de onbekende parameter(s) aan met θ , dan is dus de simultane verdelingsfunctie

$$(13.2) \quad F(w_1, \dots, w_n; \theta)$$

gegeven, op de waarde(n) van θ na. Houden de veronderstellingen in, dat de w_i onafhankelijk zijn, dan kan (13.2) als een product geschreven worden:

$$(13.3) \quad F(w_1, \dots, w_n; \theta) = \prod_{i=1}^n F_i(w_i; \theta).$$

Iedere waarde van θ komt nu overeen met een *enkelvoudige hypothese* omtrent F. Door aan θ een bepaalde waarde te geven is F immers geheel vastgelegd en bekend.*) Alle waarden, die θ aan kan nemen, vormen tezamen de *parameter-ruimte*. Ieder punt daarvan komt overeen met een hypothese omtrent F en de verzameling van al deze hypothesen wordt de verzameling der *toegelaten hypothesen* genoemd. Deze verzameling wordt dus enerzijds bepaald door de

*)

Wij zullen later gevallen van algemenere aard ontmoeten, waarbij *niet* verondersteld wordt dat F een bekende vorm bezit. De formulering wordt echterodeloos ingewikkeld indien wij deze z.g. *verdelingsvrije* gevallen er direct in betrekken.

omtrent F gemaakte veronderstellingen, anderzijds door de voor θ toegelaten (of: mogelijke) waarden.

VOORBEELD 13.1. (Vervolg)

Bij n erfelijkheidsproeven van het type van MENDEL is de verdeling van het volgende type:

$$(13.4) \quad P(w_1 = w_1 \wedge \dots \wedge w_n = w_n; p) = \prod_{i=1}^n \{w_i p + (1-w_i)q\},$$

met $w_i = 0$ of 1 ($i = 1, \dots, n$), als $p = P[A]$ de kans op verschijnsel A ($w_i = 1$) voorstelt en $q = 1 - p = P[B]$ de kans op verschijnsel B .

Deze schijnbaar ingewikkelde vorm wordt later teruggebracht tot de binomiale verdeling. Volgens de hypothese van MENDEL bezit de kans p een bepaalde waarde (waarvoor wij, om de gedachten te bepalen, de waarde $\frac{1}{4}$ aannemen), maar het doel van het statistische onderzoek zal juist zijn na te gaan, of die hypothese juist is. Worden, in principe, alle waarden van p mogelijk geacht, dan is de parameterruimte hier van de vorm

$$(13.5) \quad 0 \leq p \leq 1.$$

De te toetsen hypothese

Veelal gaat het er nu om te onderzoeken of een enkelvoudige hypothese H_0 van de vorm

$$(13.6) \quad H_0: \theta = \theta_0,$$

waarin θ_0 een gegeven getal is, op grond van de waarnemingen (13.1) al of niet verworpen moet worden. Deze hypothese wordt dan de te toetsen hypothese (ook wel: getoetste hypothese of nulhypothese) genoemd.

De te toetsen hypothese kan ook een *samengestelde hypothese* zijn, bijv. van de vorm

$$(13.7) \quad H_0: \theta \leq \theta_0 \quad \text{of} \quad \theta \geq \theta_0.$$

VOORBEELD 13.1. (Vervolg)

In het bovengenoemde geval zou de te toetsen hypothese zijn:

$$(13.8) \quad H_0: p = \frac{1}{4},$$

of in het algemeen $p = p_0$.

De alternatieve hypothesen

De uitkomst van het onderzoek is, dat de getoetste hypothese H_0 al of niet verworpen wordt. Bij verwerping van (13.6) luidt de conclusie uiteraard:

$\theta \neq \theta_0$; bij verwerping van $H_0: \theta \leq \theta_0$ luidt de conclusie: $\theta > \theta_0$.

De hypothesen, die aanvaard worden als H_0 verworpen wordt, worden de *alternatieve hypothesen* genoemd. In de regel wordt, bij verwerping van H_0 , een samengestelde (alternatieve) hypothese aanvaard.

Bij niet-verwerpen van H_0 wordt, zoals wij later zullen zien, niet tot aanvaarding van H_0 besloten. Althans dit is, strikt genomen, niet verantwoord. Niet-verwerpen betekent n.l. alleen, dat er geen voldoende aanwijzingen tegen H_0 zijn, doch dat houdt nog niet in, dat de aanwijzingen voor H_0 voldoende sterk zijn om tot aanvaarding van H_0 over te gaan. Dit punt komt nog uitvoeriger ter sprake.

VOORBEELD 13.1. (Vervolg)

Wordt (13.8) verworpen, dan luidt de conclusie:

$$(13.9) \quad p \neq \frac{1}{4},$$

In praktische bewoordingen: de erfelijkheidstheoretische hypothese, die tot de waarde $p = \frac{1}{4}$ leidde, deugt niet en zal dus door een andere (die dan gewoonlijk ingewikkelder is) vervangen moeten worden. Deze nieuwe hypothese moet dan opnieuw getoetst worden, op grond van hernieuwde waarnemingen.

De toetsingsgrootheid

De beoordeling van de te toetsen hypothese geschiedt met behulp van een *toetsingsgrootheid*, die uit de waarnemingen wordt berekend:

$$(13.10) \quad \underline{t} = t(\underline{w}_1, \dots, \underline{w}_n).$$

Deze grootheid kan meerdimensionaal zijn, maar wij zullen voor het merendeel ééndimensionale toetsingsgrootheden ontmoeten.

Gewoonlijk is $\underline{t}$ een schatter van de onbekende parameter θ , of een monotone functie daarvan (wijkt de uit de waarnemingen berekende waarde van deze schatter "sterk" van de te toetsen waarde θ_0 af, dan zal men H_0 verwerpen).

VOORBEELD 13.1. (Vervolg)

Volgens voorbeeld 12.3 (blz. 117) is

$$(13.11) \quad \hat{p} = \frac{\sum w_i}{n},$$

de meest aannemelijke schatter voor de onbekende parameter p . In plaats van $\hat{p}$ wordt $\sum w_i = n\hat{p}$ in dit geval als toetsingsgrootte gebruikt. Dit is dus: het aantal malen, dat verschijnsel A bij de n proeven waargenomen wordt.

De verdeling van $\underline{t}$ onder de getoetste hypothese

De gedachtengang van een toets komt nu ongeveer overeen met die van een bewijs uit het ongerijmde. Onder H_0 is de verdelingsfunctie F van $(w_1, \dots, w_n)$ volledig bekend en daaruit valt die van $\underline{t}$ dan af te leiden. Notatie:

$$(13.12) \quad F(t; \theta_0).$$

De verdere redenering berust grotendeels op deze verdeling.

VOORBEELD 13.1. (Vervolg)

Stellen wij het aantal malen A voor door:

$$(13.13) \quad \underline{x} = \sum w_i,$$

dan bezit $\underline{x}$, onder $H_0 : p = \frac{1}{4}$, een binomiale verdeling:

$$(13.14) \quad P(\underline{x}=x; \tfrac{1}{4}) = \binom{n}{x} \left(\tfrac{1}{4}\right)^x \left(\tfrac{3}{4}\right)^{n-x} \quad (x = 0, \dots, n).$$

De kritieke zone

Onder H_0 kan $\underline{t}$ allerlei waarden aannemen. Meestal zijn dit dezelfde waarden als die onder alternatieve hypothesen aangenomen kunnen worden (alleen met andere kansen). Het gebied, dat $\underline{t}$ kan doorlopen, geven wij aan met T :

$$(13.15) \quad P(\underline{t} \in T; \theta_0) = 1 \text{ voor alle } \theta_0 \text{ gespecificeerd door } H_0.$$

Een toets komt nu daarop neer, dat T in twee of meer gebieden verdeeld wordt en dat de te trekken conclusie afhangt van het gebied, waarin de door $\underline{t}$ aangenomen waarde t terecht komt.

Deze gebieden hebben gewoonlijk de volgende vorm:

$$(13.16) \quad t \leq t_{\ell} \quad \text{leidt tot de conclusie: } \theta < \theta_0,$$

$$(13.17) \quad t \geq t_r \quad " \quad " \quad " \quad " \quad : \theta > \theta_0,$$

$$(13.18) \quad t_{\ell} < t < t_r \quad " \quad " \quad " \quad " \quad : \theta = \theta_0 \text{ wordt niet verworpen.}$$

Hierbij is verondersteld, dat t een schatter van θ is, zodat een grote (kleine) waarde van t een aanwijzing vormt dat θ een grote (kleine) waarde bezit; is de toetsingsgrootheid een monotoon stijgende functie van een schatter van θ , dan blijft de situatie onveranderd; is hij een monotoon dalende functie daarvan, dan verwisselen de conclusies van (13.16) en (13.17). De keuze van t_{ℓ} en t_r komt later ter sprake.

Het schema (13.16) - (13.18) heeft betrekking op een *tweezijdige toets*, waarbij een enkelvoudige hypothese van de vorm $\theta = \theta_0$ (of een samengestelde van de vorm $\theta'_0 \leq \theta \leq \theta_0$) getoetst wordt tegen de alternatieve hypothese $\theta < \theta_0$ (resp. $\theta < \theta'_0$) en $\theta > \theta_0$. De gebieden $t \leq t_{\ell}$ resp. $t \geq t_r$ worden ook met Z_{ℓ} resp. Z_r aangegeven en heten de *linker-* resp. *rechterkritieke zone*; zij worden ook tezamen de *(tweezijdige) kritieke zone* genoemd.

Is de getoetste hypothese van de vorm $\theta \geq \theta_0$ (resp. $\theta \leq \theta_0$) dan is de alternatieve hypothese *éénzijdig*, n.l. $\theta < \theta_0$ (resp. $\theta > \theta_0$) en dan wordt het schema als volgt:

$$(13.19) \quad t \leq t_{\ell} \text{ leidt tot de conclusie: } \theta < \theta_0,$$

$$(13.20) \quad t > t_{\ell} \quad " \quad " \quad " \quad " \quad : \theta \geq \theta_0 \text{ wordt niet verworpen,}$$

(*linkséénzijdige toets*), resp.

$$(13.21) \quad t \geq t_r \text{ leidt tot de conclusie: } \theta > \theta_0,$$

$$(13.22) \quad t < t_r \quad " \quad " \quad " \quad " \quad : \theta \leq \theta_0 \text{ wordt niet verworpen,}$$

(*rechtséénzijdige toets*).

De getallen t_{ℓ} en t_r heten de *linker-* resp. *rechter kritieke waarden*.

VOORBEELD 13.1. (Vervolg)

De getoetste hypothese is: $p = \frac{1}{4}$. Worden nu zeer weinig proeven met A als

uitkomst verkregen (dus kleine x), dan wordt $p = \frac{1}{4}$ verworpen ten gunste van $p < \frac{1}{4}$; worden er echter zeer veel gevonden, dan concludeert men $p > \frac{1}{4}$. Wij hebben hier dus een geval van een tweezijdige toets. De keuze der kritieke waarden komt later. Het is echter wel reeds intuïtief duidelijk, dat het dwaasheid zou zijn de hypothese $p = \frac{1}{4}$, bij bijv. $n = 1000$, wel te verwerpen als $x = 780$ is, maar niet als $x > 780$ is. Dit leidt dus inderdaad tot kritieke zones "in de staarten" van het gebied, dat x kan doorlopen.

De onbetrouwbaarheid van een toets

Indien men op deze wijze te werk gaat is het klaarblijkelijk niet uitgesloten, dat H_0 verworpen wordt, als deze hypothese juist is. Deze foute uitkomst (een juiste hypothese verwerpen) wordt een *fout van de eerste soort* genoemd. Er is slechts één methode om deze fout met zekerheid uit te sluiten en dat is: H_0 nooit verwerpen. Immers, zodra men een niet-lege kritieke zone gebruikt is er, ook als H_0 juist is, een positieve kans H_0 te verwerpen.

Deze kans op een fout van de eerste soort wordt kortweg de *onbetrouwbaarheid* α_0 van de toets genoemd. Deze α_0 mag uiteraard niet groot zijn en men eist daarom, dat de kritieke zone zo gekozen wordt, dat voldaan is aan

$$(13.23) \quad \alpha_0 \leq \alpha,$$

waarin α een door de onderzoeker gekozen waarde bezit.

De grootte α wordt de *onbetrouwbaarheidsdrempel* genoemd. Hiervoor wordt veelal de waarde 0,05 genomen, zonder dat daarvoor een andere fundering bestaat dan de traditie. Indien een fout van de eerste soort ernstige gevolgen heeft kiest men ook wel kleinere waarden, bijv. 0,01 of 0,001.

Het onderscheid tussen α_0 en α lijkt wellicht overbodig. Waarom niet $\alpha_0 = \alpha$ gemaakt? Inderdaad; dit doet men steeds als dat mogelijk is. Maar als de toetsingsgrootte t onder H_0 een discrete verdeling bezit is dit in de regel niet mogelijk. Men neemt dan α_0 zo dicht bij α als deze verdeling en (13.23) toestaan.

VOORBEELD 13.1. (Vervolg)

Als $n = 50$ is en $H_0: p = \frac{1}{4}$, dan ziet de verdelingsfunctie $F(x) = P(\underline{x} \leq x)$ er - in 3 decimalen berekend - volgens (13.14) uit als in tabel 13.1. In deze tabel is, behalve $F(x)$, ook opgenomen $R(x) = P(\underline{x} \geq x) = 1 - F(x-1)$.

Nemen wij nu voor Z_L , de linkerkritieke zone, $x \leq 6$ en voor Z_R , de rechterkritieke zone, $x \geq 19$, dan is volgens deze tabel:

Tabel 13.1F(x) en R(x) van de binomiale verdeling met $n = 50$, $p = \frac{1}{4}$

x	F(x)	R(x)
0	0,000001	1
1	0,00001	≈ 1
2	0,0001	≈ 1
3	0,0005	0,9999
4	0,002	0,9995
5	0,007	0,998
6	0,019	0,993
7	0,045	0,981
8	0,092	0,955
9	0,164	0,918
10	0,262	0,836
11	0,382	0,738
12	0,510	0,618
13	0,637	0,490
14	0,748	0,363
15	0,837	0,252
16	0,902	0,163
17	0,945	0,098
18	0,971	0,055
19	0,986	0,029
20	0,994	0,014
21	0,997	0,006
22	0,999	0,003
23	0,9996	0,001
24	0,9998	0,0004
25	0,9999	0,0002
26,...,50	≈ 1	≈ 0

(Deze tabel is ontleend aan H.G. ROMIG, *50-100 Binomial Tables*, Wiley, N.Y., Chapman and Hall, London, 1953.)

$$\alpha_0 = \alpha_l + \alpha_r = P(\underline{x} \leq 6; \frac{1}{4}) + P(\underline{x} \geq 19; \frac{1}{4}) =$$

$$= 0,019 + 0,029 = 0,048.$$

Is de onbetrouwbaarheidsdrempel $\alpha = 0,05$ en wenst men de beide delen van de kritieke zone ongeveer symmetrisch te behandelen - d.w.z. allebei ongeveer $\frac{1}{2}\alpha = 0,025$ tot de onbetrouwbaarheid te laten bijdragen - dan is dit

de meest geschikte kritieke zone. Wij zien, dat $\alpha_0 < \alpha$ is.

Voor $\alpha = 0,10$ wordt $Z_\ell: x \leq 7$ en $Z_r: x \geq 18$. In dit geval is (toevalligerwijze) $\alpha_0 = \alpha$.

Voor $\alpha = 0,01$ is $Z_\ell: x \leq 5$ en $Z_r: x \geq 22$ met $\alpha_0 = \alpha = 0,01$; een andere mogelijkheid is:

$$Z_\ell: x \leq 4 \quad \text{en} \quad Z_r: x \geq 21, \quad \text{met} \quad \alpha_0 = 0,008.$$

Een algemeen aanvaarde regel, om in dergelijke gevallen tussen de verschillende mogelijkheden te kiezen, is niet beschikbaar.

Linker- en rechter-onbetrouwbaarheid

Zoals uit de behandeling van het voorbeeld al blijkt, bestaat de onbetrouwbaarheid van een tweezijdige toets voor de hypothese $H_0: \theta = \theta_0$ uit twee delen, behorende bij het linker- en rechterdeel van de kritieke zone afzonderlijk. Wij hebben:

$$(13.24) \quad \alpha_0 = \alpha_\ell + \alpha_r,$$

met

$$(13.25) \quad \alpha_\ell = P(\underline{t} \leq t_\ell; \theta_0) \quad (\text{linkeronbetrouwbaarheid})$$

en

$$(13.26) \quad \alpha_r = P(\underline{t} \geq t_r; \theta_0) \quad (\text{rechteronbetrouwbaarheid}).$$

Men noemt een toets *symmetrisch tweezijdig* als voor t_ℓ en t_r de grootste resp. kleinste waarde van t genomen wordt, waarvoor α_ℓ en α_r beide $\leq \frac{1}{2}\alpha$ zijn.

De in het voorbeeld genoemde tweezijdige kritieke zones zijn dus niet geheel symmetrisch, maar komen daar wel bij in de buurt.

Als $\underline{t}$ een continue verdeling bezit (hetgeen bij het beschouwde voorbeeld niet het geval is), werkt men in het tweezijdige geval gewoonlijk met symmetrische kritieke zones. Ook als men benaderingsmethoden gebruikt wordt dit gewoonlijk gedaan, zoals verderop in dit hoofdstuk ook zal geschieden.

Het onderscheidingsvermogen (Eng.: *power function*)

Het ten onrechte verwerpen van een (juiste) hypothese - een fout van de eerste soort - vindt zijn tegenhanger in het ten onrechte *niet* verwerpen

van een (onjuiste) hypothese - een fout van de tweede soort.

Een fout van de eerste soort kan alleen gemaakt worden als H_0 : $\theta = \theta_0$ juist is en de kans daarop is in dat geval gelijk aan α_0 .

Een fout van de tweede soort is alleen mogelijk als H_0 onjuist is. In dat geval is een alternatieve hypothese H : $\theta = \theta'$ met $\theta' \neq \theta_0$ juist en de kans op het maken van een fout van de tweede soort is afhankelijk van deze θ' . Deze kans is:

$$(13.27) \quad 1 - \pi(\theta') = 1 - P(t \in Z; \theta'), \quad \text{voor elke } \theta' \text{ gespecificeerd door } H,$$

waarin Z de kritieke zone voorstelt. Immers als t in Z ligt, wordt H_0 (niet H) verworpen en dan wordt dus geen fout van de tweede soort gemaakt. De grootte $\pi(\theta)$ wordt het *onderscheidingsvermogen* van de toets genoemd. Dit is een functie van $\theta \in \Theta$ waarvoor geldt:

$$(13.28) \quad \pi(\theta) = \alpha_0, \quad \text{voor alle } \theta \text{ gespecificeerd door } H_0.$$

Voor $\theta' \neq \theta_0$ is $\pi(\theta')$ de kans op verwerping van H_0 : $\theta = \theta_0$ als H : $\theta = \theta'$ juist is en H_0 dus ook verworpen dient te worden. Dit is dus, om zo te zeggen, de kans, dat de toets H van H_0 zal weten te onderscheiden.

Bij een tweezijdige toets zijn drie uitkomsten mogelijk, n.l. - als θ de onbekende parameter is en θ_0 de getoetste waarde - de drie uitspraken:

$$\theta < \theta_0,$$

$$\theta > \theta_0$$

en

$$\theta = \theta_0 \quad \text{wordt niet verworpen.}$$

De kansen op deze drie uitspraken geven wij aan met $\pi_l(\theta)$, $\pi_r(\theta)$ en $1 - \pi(\theta)$, zodat dus

$$(13.29) \quad \pi_l(\theta) + \pi_r(\theta) = \pi(\theta)$$

is. De beide termen heten het *linker-* resp. *rechteronderscheidingsvermogen*.

VOORBEELD 13.1. (Vervolg)

Met behulp van een tabel van de binomiale verdeling kan men nu gemakkelijk het onderscheidingsvermogen van een binomiale toets - als functie van de onbekende kans p - bepalen. Wij doen dit voor de kritieke zones

$$\text{en} \quad \left. \begin{array}{l} Z_l: x \leq 6 \quad \text{met } \alpha_l = 0,019 \\ Z_r: x \geq 19 \quad \text{met } \alpha_r = 0,029 \end{array} \right\} n = 50; p_0 = \frac{1}{4}.$$

In de reeds genoemde tabel van ROMIG kunnen wij nu voor verschillende waarden van p de waarde van

$$(13.30) \quad \pi_{\ell}(p) = P(\underline{x} \leq 6; p)$$

en van

$$(13.31) \quad \pi_r(p) = P(\underline{x} \geq 19; p) = 1 - P(\underline{x} \leq 18; p)$$

opzoeken. Deze waarden staan in tabel 13.2 vermeld, evenals hun som: $\pi(p)$.

In fig. 13.1 is deze tabel grafisch in beeld gebracht.

Tabel 13.2

Het onderscheidingsvermogen van de binomiale toets met
 $n = 50$ en $Z_{\ell}: x \leq 6$, $Z_r: x \geq 19$

p	$\pi_{\ell}(p)$	$\pi_r(p)$	$\pi(p)$
0,05	0,988	0,000	0,988
0,10	0,770	0,000	0,770
0,15	0,361	0,000	0,362
0,20	0,103	0,003	0,106
0,21	0,076	0,004	0,080
0,22	0,056	0,008	0,064
0,23	0,040	0,012	0,052
0,24	0,028	0,019	0,047
0,25	0,019	0,029	0,048
0,26	0,013	0,042	0,055
0,30	0,002	0,141	0,143
0,40	0,000	0,664	0,664
0,50	0,000	0,968	0,968

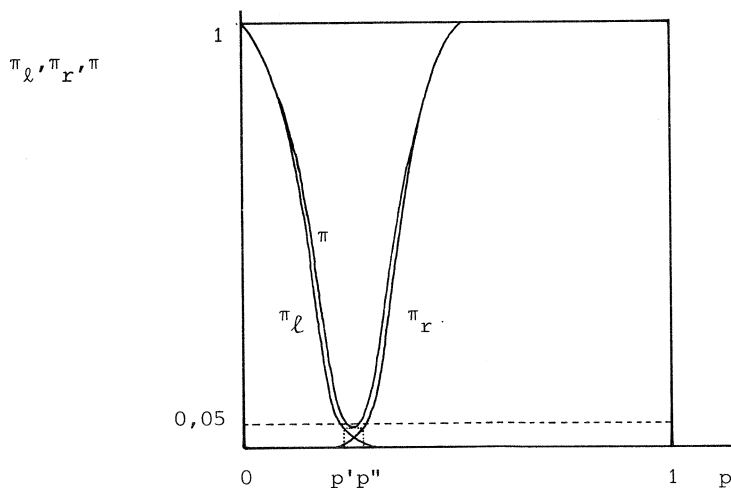


Fig.13.1. Het onderscheidingsvermogen van de binomiale toets met

$n = 50$, $Z_{\ell}: x \leq 6$ en $Z_r: x \geq 19$.

Bij tabel 13.2 en fig. 13.1 zijn nog verschillende OPMERKINGEN te maken.

1. Hoe steiler het onderscheidingsvermogen verloopt, hoe beter de toets is.

Dit volgt direct uit de betekenis van deze functie: de kans op verworping van H_0 , Indien H_0 onjuist is. Men kan nu bewijzen - hetgeen wij hier nalaten - dat voor toenemende n , bij constante α , de beide krommen π_ℓ en π_r tot limiet een verticale lijn door $p = \frac{1}{4}$ als abscis hebben. Wordt een andere waarde, $p = p_0$, getoetst, dan worden daarvoor uiteraard andere kritieke zones gebruikt en dan geldt deze stelling analoog, nu met een verticale lijn door $p = p_0$ als limiet.

2. De toets, die hier geschreven is als een toets voor de hypothese $p = \frac{1}{4}$ kan evengoed (of zelfs beter) beschouwd worden als een toets voor $p = 0,24$, met $\alpha_0 = 0,047$. In feite kan men, met $\alpha = 0,05$, zeggen dat het een toets is voor de samengestelde hypothese $p' \leq p \leq p''$, waarin p' en p'' de abscis voorstellen van de snijpunten van de horizontale rechte $\pi = 0,05$ met de kromme $\pi(p)$. Want voor alle tussenliggende waarden voor p is de kans dat $\underline{x}$ in Z_ℓ of Z_r valt $\leq \alpha$. Deze fijne onderscheidingen worden gewoonlijk niet gemaakt.

3. De kromme π_ℓ stelt de kans voor, dat tot $p < 0,25$ besloten wordt. Deze kans is ook > 0 als p in werkelijkheid $> 0,25$ is, maar hij neemt (gelukkig) met toenemende p sterk af. Voor $p = 0,30$ stelt nu $\pi_\ell(0,30) = 0,002$ de kans voor om te besluiten tot $p < 0,25$, een falikant verkeerde conclusie; $\pi_r(0,30) = 0,141$ is daarentegen juist de kans op de conclusie $p > 0,25$, die dan juist is. De som $\pi(0,30)$ van $\pi_\ell(0,30)$ en $\pi_r(0,30)$ is dus eigenlijk de som van de kansen op twee volkomen ongelijksoortige conclusies, n.l. een falikant onjuiste en een juiste. Vanuit het standpunt van de toetsingstheorie bezien is het daarom beter de merites van een toets te beoordelen op grond van π_ℓ en π_r afzonderlijk dan op grond van hun som (vgl. echter stelling 15.2).

4. Is de toets éézijdig, dan ontbreekt één van de twee krommen π_ℓ of π_r .

5. Men kan bewijzen, dat bij gegeven n , p_0 en α het grootste onderscheidingsvermogen verkregen wordt als de kritieke zones "in de staarten" genomen worden, zoals wij hier gedaan hebben. Zou men een kritieke zone nemen van het type:

$$Z: 0 < x_1 \leq x \leq x_2 < n$$

dan verkrijgt men een toets met een zeer ongunstig onderscheidingsvermogen, zoals wij later met een voorbeeld zullen laten zien.

Zuiverheid en asymptotische onderscheidenheid

DEFINITIE 13.1. Een toets wordt zuiver genoemd, als het onderscheidingsvermogen voor alle alternatieve hypothesen groter is dan de onbetrouwbaarheid.

De praktische betekenis is, dat bij een zuivere toets de kans om de getoetste hypothese te verwerpen als deze juist is kleiner is dan wanneer deze onjuist is.

DEFINITIE 13.2. Een toets voor de hypothese $H_0: \theta = \theta_0$ heet *asymptotisch onderscheidend* met betrekking tot een alternatieve hypothese $H: \theta = \theta_1$, als geldt:

$$(13.32) \quad \lim_{n \rightarrow \infty} \pi(\theta_1) = 1.$$

De praktische betekenis hiervan is, dat men, als H juist is, bij een voldoende groot aantal waarnemingen vrijwel zeker is dat de toets (terecht) tot verwerping van H_0 zal leiden.

Eénzijdige toetsen zijn vrijwel altijd zuiver. Het onderscheidend vermogen is dan n.l. meestal een monotone kromme, zoals in fig. 13.1 (π_L en π_R afzonderlijk). Tweezijdige toetsen zijn lang niet altijd zuiver, ook niet als dit voor de beide éénzijdige wel geldt.

Toetsen, die voor relevante alternatieve hypothesen niet asymptotisch onderscheidend zijn, zijn voor de praktijk in de regel onbruikbaar.

VOORBEELD 13.1. (Vervolg)

Uit fig. 13.1 blijkt, dat in dit geval de tweezijdige toets, beschouwd als toets voor de hypothese $p = \frac{1}{4}$, niet zuiver is: immers voor $p = 0,24$ is het onderscheidingsvermogen kleiner dan de onbetrouwbaarheid (vgl. tabel 13.2). Beschouwen wij echter $H_0: p' \leq p \leq p''$ als de getoetste hypothese, dan is de toets wel zuiver.

De onzuiverheid ontstaat door 2 oorzaken:

1. het discrete karakter van de verdeling van de toetsingsgrootte $\underline{x}$, waardoor men de beide delen van de kritieke zone niet volledig "tegen elkaar afwegen" kan;
2. doordat de verdeling van $\underline{x}$, onder $H_0: p = \frac{1}{4}$, asymmetrisch is, hetgeen het zuiver maken van de toets weliswaar niet onmogelijk maakt, maar wel bemoeilijkt.

Heeft de toetsingsgrootheid onder H_0 een symmetrische verdeling, dan behoeft men de tweezijdige kritieke zone slechts symmetrisch te nemen om zuiverheid te bereiken ook voor de tweezijdige toets.

De toets is asymptotisch onderscheidend voor alle van $\frac{1}{2}$ verschillende waarden van p ; dit is boven (onder OPMERKINGEN, 1.) reeds - zonder bewijs - vermeld.

Overschrijdingskans

De bij een waarde t van de toetsingsgrootheid behorende *linker* resp. *rechteroverschrijdingskans* is:

$$(13.33) \quad k_l(t) \stackrel{\text{def}}{=} P(\underline{t} \leq t; \theta_0) \quad \text{resp.} \quad k_r(t) \stackrel{\text{def}}{=} P(t \geq \underline{t}; \theta_0).$$

Is

$$(13.34) \quad k_l(t) \leq \alpha_l \quad \text{resp.} \quad k_r(t) \leq \alpha_r,$$

dan ligt t in Z_l resp. Z_r .

De *tweezijdige overschrijdingskans* $k(t)$, behorende bij een symmetrische tweezijdige toets, wordt gedefinieerd als het dubbele van de kleinste der beide eenzijdige:

$$(13.35) \quad k(t) \stackrel{\text{def}}{=} 2 \min\{k_l(t), k_r(t)\}.$$

Dan geldt: is

$$(13.36) \quad k(t) \leq \alpha,$$

dan is $t \in Z$. Voor een niet-symmetrische wordt de tweezijdige overschrijdingskans op iets andere wijze gedefinieerd.

VOORBEELD 13.1. (Vervolg)

Volgens tabel 13.1 is $k_l(3) = 0,0005$, $k_r(21) = 0,006$, $k(4) = 0,004$.

De betekenis van het begrip overschrijdingskans is, dat men aan de grootte ervan niet alleen kan zien of de beschouwde waarde van t in de kritieke zone van de toets ligt, maar ook of dit maar juist, of ruim het geval is. Hoe kleiner de gevonden overschrijdingskans is, des te stelliger kan men de getoetste hypothese verwerpen.

De belangrijkste elementaire begrippen van de toetsingstheorie zijn

hiermede in het kort behandeld. Zij vormen tezamen weliswaar slechts een begin van de toetsingstheorie, maar zij zijn voldoende om er een aantal toetsen op te baseren. Wij behandelen verder in dit hoofdstuk alleen de in de praktijk veelvuldig gebruikte techniek van de binomiale toetsen met behulp van de normale verdeling als benadering van de binomiale.

Binomiale toetsen met normale benadering

In hoofdstuk XI (stelling 11.8 en figuur 11.1) hebben wij gezien, dat de binomiale verdeling met behulp van de aangepaste normale verdeling benaderd kan worden. Wij brengen nu op deze benadering nog een verfijning aan, die bekend staat onder de naam *continuïteitscorrectie*. Deze correctie heeft betrekking op het feit, dat de *continue* normale verdeling gebruikt wordt om de *discrete* binomiale te benaderen.

Indien n en p de parameters van de binomiale verdeling zijn, zodat

$$(13.37) \quad P(\underline{x}=x) = \binom{n}{x} p^x q^{n-x}$$

is, met $E\underline{x} = np$ en $\sigma^2(\underline{x}) = npq$, dan wordt de normale verdeling met dezelfde verwachting en variantie voor de benadering gebruikt.

Een discrete kans $P(\underline{x}=x)$ moet nu benaderd worden door een oppervlak onder de normale verdelingsdichtheid en het ligt voor de hand hiervoor het oppervlak boven het interval $(x-\frac{1}{2}, x+\frac{1}{2})$ te nemen, dus symmetrisch om de waarde x gelegen (vgl. fig. 13.2). Alleen voor $P(\underline{x}=0)$ en voor $P(\underline{x}=n)$ neemt men de staart er volledig bij, dus dan neemt men het "normale" oppervlak boven het interval $(-\infty, \frac{1}{2})$ resp. $(n-\frac{1}{2}, \infty)$, omdat anders de som der benaderende kansen kleiner dan 1 zou zijn.

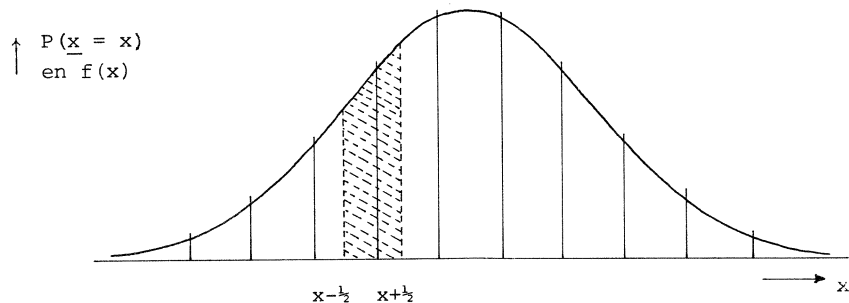


Fig. 13.2. Normale benadering van een binomiale verdeling met continuïteitscorrectie.

Wenst men nu de *linkeroverschrijdingskans* $P(\underline{x} \leq x)$ te benaderen, dan neemt men daarvoor dus het "normale" oppervlak boven het interval $(-\infty, x + \frac{1}{2})$, voor de *rechteroverschrijdingskans* $P(\underline{x} \geq x)$ dat boven het interval $(x - \frac{1}{2}, \infty)$. Deze benaderende oppervlakken kunnen, na standaardisering, uit tabel 1 van de normale verdeling afgelezen worden. De term $\frac{1}{2}$ is de continuïteitscorrectie.

In formule wordt dit:

$$(13.38) \quad k_l(x) = P(\underline{x} \leq x; n, p) \approx P\left(\underline{u} \leq \frac{x - np + \frac{1}{2}}{\sqrt{npq}}\right),$$

resp.

$$(13.39) \quad k_r(x) = P(\underline{x} \geq x; n, p) \approx P\left(\underline{u} \geq \frac{x - np - \frac{1}{2}}{\sqrt{npq}}\right),$$

en voor de *tweezijdige overschrijdingskans*:

$$(13.40) \quad k(x) \approx 2P\left(\underline{u} \geq \frac{|x - np| - \frac{1}{2}}{\sqrt{npq}}\right).$$

VOORBEELD 13.1. (Vervolg)

Voor $n = 50$ en $p = \frac{1}{4}$ geldt volgens tabel 13.1

$$k_r(18) = P(\underline{x} \geq 18; n = 50, p = \frac{1}{4}) = 0,055.$$

Wij passen nu de normale benadering met continuïteitscorrectie toe om deze zelfde overschrijdingskans te benaderen: $Ex = np = 12,5$, $\sigma^2(\underline{x}) = npq = 9,375$, dus $\sigma(\underline{x}) = 3,06$.

Derhalve

$$\begin{aligned} P(\underline{x} \geq 18) &= P\left(\frac{x - 12,5}{3,06} \geq \frac{18 - 12,5}{3,06}\right) \approx P\left(\underline{u} \geq \frac{17,5 - 12,5}{3,06}\right) = \\ &= P(\underline{u} \geq 1,63) = 0,0516. \end{aligned}$$

Zonder continuïteitscorrectie zouden wij vinden

$$P(\underline{x} \geq 18) \approx P\left(\underline{u} \geq \frac{18 - 12,5}{3,06}\right) = P(\underline{u} \geq 1,79) = 0,0367.$$

De benadering met continuïteitscorrectie is beter, hetgeen echter niet altijd het geval is. Wel geeft de benadering met continuïteitscorrectie

altijd een grotere uitkomst voor een overschrijdingskans dan zonder, zodat minder gauw tot verwerping van de getoetste hypothese wordt overgegaan. In het onderhavige geval zou bij rechtséénzijdige toetsing met onbetrouwbaarheidsdrempel 0,05 bij gebruik van de exacte overschrijdingskans of de benadering met continuïteitscorrectie op grond van $x = 18$ niet tot verwerping van $H_0: p = \frac{1}{4}$ worden overgegaan, maar zonder continuïteitscorrectie wel. Het is daarom "voorzichtiger" de continuïteitscorrectie te gebruiken dan deze weg te laten. Voor grote n is het effect echter gering, zodat weglating dan wel gerechtvaardigd geacht mag worden.

OPGAVE 13.1. Bereken voor $n = 50$, $p = \frac{1}{4}$, de linkeroverschrijdingskans van $x = 6$; exact, met en zonder continuïteitscorrectie.
(Antwoord: 0,019; 0,025; 0,017.)

OPGAVE 13.2. Bereken voor $n = 900$ en $p = 0,5$ de tweezijdige overschrijdingskans van de waarde $x = 427$ met en zonder continuïteitscorrectie.

OPGAVE 13.3. Bij 225 onafhankelijke experimenten, die alle dezelfde kans p op succes hebben worden 85 successen gevonden.

a) Toets de hypothese, dat $p = \frac{1}{2}$ is, met onbetrouwbaarheidsdrempel 0,05.

b) Toets, met dezelfde onbetrouwbaarheidsdrempel, de hypothese dat $p = \frac{1}{3}$.
Voer beide toetsen tweezijdig en linkséénzijdig uit.

(Opgave 5 van het Sommenboekje; oplossing blz. 50-51.)

OPGAVE 13.4. Van een binomiale verdeling is $n = 10$ en $p = 0,3$.

a) Bereken met behulp van de normale benadering van de binomiale verdeling de linkeroverschrijdingskans van $x = 1$, met en zonder continuïteitscorrectie.

b) Bereken deze kans ook exact en bereken vervolgens hoe groot de continuïteitscorrectie in dit geval genomen zou moeten worden, om te bereiken dat de benaderende en de exacte overschrijdingskans precies gelijk zijn.

(Opgave 26 van het Sommenboekje; oplossing blz. 71.)

OPGAVE 13.5. Een organisator van een vuurwerk weet uit ervaring, dat van de vuurpijlen, die hij gewoonlijk gebruikt, ongeveer één op de 5 weigert. In verband daarmee maakt hij een aantal reserve pijlen gereed. Hij wenst het publiek minstens 20 goede vuurpijlen te laten zien en hij aanvaardt

hoogstens een kans van $1/100$, dat hem dit niet zal lukken. Hoeveel reserve pijlen moet hij opstellen?

(Opgave 40 van het Sommenboekje; oplossing blz. 83.)

OPGAVE 13.6. Wat is, voor $n = 50$ en $H_0: p = \frac{1}{4}$, de onbetrouwbaarheid van de kritieke zone, die uit de waarden 18, 19 en 20 bestaat?

Bereken het onderscheidingsvermogen van deze kritieke zone voor $p = 0,2; 0,4; 0,6; 0,8$ en 1 met behulp van de normale benadering en vergelijk dit met π_r in fig. 13.1 en tabel 13.2.

(Antwoord:

$$\alpha_0 = 0,049,$$

p	0,02	0,25	0,4	0,6	0,8	1
$\pi \approx$	0,007	0,049	0,37	0,001	0	0 .)

Ook voor de *kritieke waarden* zijn nu eenvoudige benaderingsformules af te leiden. Is u_α de waarde van een $N(0,1)$ -grootheid, waarvoor

$$(13.41) \quad P(\underline{u} \geq u_\alpha) = \alpha$$

is $-u_\alpha$ is dus in tabel 1 te vinden voor $k_r = \alpha$ - dan zijn de kritieke waarden van de linker- en rechterzijdige toets

$$(13.42) \quad x_{\ell} \approx np_0 - \left\{ u_{\alpha_{\ell}} \sqrt{np_0 q_0} + \frac{1}{2} \right\},$$

respectievelijk

$$(13.43) \quad x_r \approx np_0 + \left\{ u_{\alpha_r} \sqrt{np_0 q_0} + \frac{1}{2} \right\},$$

terwijl voor de symmetrische tweezijdige toets geldt:

$$(13.44) \quad x_{r,\ell} \approx np_0 \pm \left\{ u_{\frac{1}{2}\alpha} \sqrt{np_0 q_0} + \frac{1}{2} \right\}.$$

De termen $\frac{1}{2}$ stellen de continuïteitscorrectie voor.

OPGAVE 13.7. Geef de afleiding van deze formules.

OPGAVE 13.8. Hoe groot moet men het aantal experimenten minstens nemen om

bij rechtséénzijdige toetsing van de hypothese $p = 0,3$ met $\alpha_r = 0,05$ voor $p = 0,6$ een onderscheidingsvermogen van mintens 0,9 te bereiken? (De continuïteitscorrectie mag verwaarloosd worden. Houd bij de oplossing rekening met het feit, dat de toetsingsgrootte alleen gehele waarden aan kan nemen.)

(Opgave 58 van het Sommenboekje; oplossing blz. 97-98.)

De keuze tussen één- of tweezijdig toetsen hangt af van het doel van het onderzoek en van de toegelaten hypothesen. Wij zullen hierover kort zijn.

Indien men een beslissing moet nemen tussen twee handelwijzen, bijv. het al of niet aanschaffen van een nieuwe machine, dan zal men daartoe alleen wensen te besluiten als de nieuwe machine beter is dan de oude. Is nu van de oude machine bekend, dat deze een fractie p_0 aan uitval geeft, dan toetst men de hypothese $H_0: p \geq p_0$ (waarin p de fractie uitval van de nieuwe machine voorstelt) tegen de alternatieve hypothese $p < p_0$. Alleen bij verwerping van H_0 gaat men tot aanschaf van de nieuwe machine over. De onbetrouwbaarheid α_ℓ van de toets stelt dan de kans voor, dat men de nieuwe machine toch zal aanschaffen indien de oude even goed is. Is de oude beter, dan is de kans op deze foute beslissing kleiner dan α_ℓ . Een tweezijdige toets zou hier niet op zijn plaats zijn, daar verwerping van $H_0: p = p_0$ ten gunste van $p > p_0$ evenmin tot aanschaffing van de nieuwe machine leidt als niet-verwerpen van H_0 . En het onderscheidingsvermogen van deze rechtséénzijdige toets is, bij dezelfde onbetrouwbaarheidsdrempel, voor $p < p_0$ het grootst voor de rechtséénzijdige toets, zodat deze een kleinere kans dan de tweezijdige geeft om een fout van de tweede soort te maken (hetgeen hier zou zijn: de nieuwe machine niet aanschaffen, hoewel hij beter is dan de oude).

Het kan ook voorkomen, dat hypothesen $p > p_0$ onmogelijk geacht worden. Deze situatie doet zich bijv. voor bij het toedienen van middelen, waarvan men met recht kan zeggen: "baat het niet, dan schaadt het toch in ieder geval niet". In zo'n geval is het gebruik van een tweezijdige kritieke zone zinloos, daar het rechterdeel hiervan behoort bij de dan zeker onjuiste conclusie, dat $p > p_0$ is.

Bij de meeste wetenschappelijke onderzoeken zal men echter met beide mogelijkheden ($p < p_0$ en $p > p_0$) rekening moeten houden en dus tweezijdig moeten toetsen.

ENIGE GANGBARE TOETSEN

14.1. De 2×2 -tabel (dubbele dichotomie)

In het vorige hoofdstuk zijn toetsen behandeld, die op de binomiale verdeling berusten en waarmee de hypothese $H_0: p = p_0$ getoets kan worden. Wij behandelen nu een toets voor de hypothese

$$(14.1.1) \quad H_0: p_1 = p_2,$$

inhoudende, dat twee onbekende kansen gelijk zijn.

Dit is een samengestelde hypothese. De parameterruimte (zie fig. 14.1.1) is een deel van het (p_1, p_2) -vlak begrensd door de ongelijkheden $0 \leq p_1 \leq 1$; $0 \leq p_2 \leq 1$ en H_0 wordt voorgesteld door de diagonaal van dat gebied.

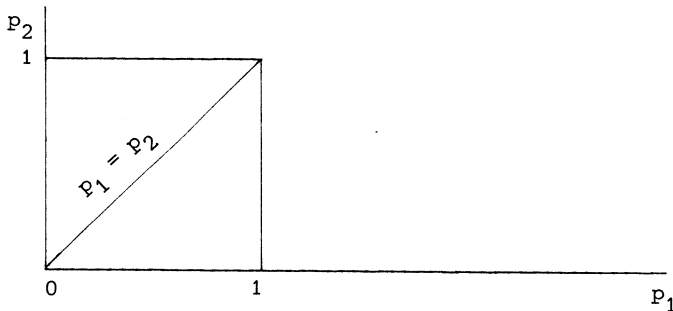


Fig. 14.1.1. De parameterruimte en H_0 van de toets voor gelijkheid van twee onbekende kansen.

De toets, die wij gaan beschrijven, kan toegepast worden indien het waarnemingsmateriaal bestaat uit twee onafhankelijke reeksen van onafhankelijke experimenten, waarbij p_1 resp. p_2 de kansen op "succes" bij de

experimenten van de eerste resp. tweede reeks voorstellen. Zijn de aantallen experimenten n resp. m en de aantallen successen $\underline{a}$ resp. $\underline{b}$, dan kunnen deze resultaten samengevat worden in een tabel van de vorm van tabel 14.1.1.

Tabel 14.1.1.
 2×2 -tabel voor $H_0: p_1 = p_2$

	Aantal malen		
	succes	mislukking	totaal
eerste reeks	$\underline{a}$	$\underline{c}$	n
tweede reeks	$\underline{b}$	$\underline{d}$	m
totaal	$\underline{r}$	$\underline{s}$	N

Hierin is:

$$(14.1.2) \quad \underline{a} + \underline{c} = n; \quad \underline{b} + \underline{d} = m; \quad \underline{r} = \underline{a} + \underline{b}; \quad \underline{s} = \underline{c} + \underline{d}; \quad N = n + m = \underline{r} + \underline{s}.$$

Een dergelijke tabel wordt een 2×2 -tabel of ook een *dubbele dichotomie* genoemd.

Als *toetsingsgrootte* wordt de stochastische grootte $\underline{a}$ genomen. De verdeling van $\underline{a}$, onder H_0 , is echter niet ondubbelzinnig bepaald, maar hangt nog af van de onbekende waarde $p_1 (= p_2)$. Deze onbekende parameter kan echter in dit geval geëlimineerd worden door niet met de onvoorwaardelijke verdeling van $\underline{a}$ (onder H_0) te werken, maar met de *voorwaardelijke verdeling* van $\underline{a}$ onder de voorwaarde $\underline{r} = r$ (hetgeen impliceert $\underline{s} = N - r$). Deze voorwaardelijke verdeling is, als $H_0: p_1 = p_2$ waar is, de hypergeometrische:

$$(14.1.3) \quad P(\underline{a} = a | \underline{r} = r) = \frac{\binom{n}{a} \binom{m}{b}}{\binom{N}{r}},$$

of ook

$$(14.1.4) \quad P(\underline{a} = a | \underline{r} = r) = \frac{\binom{r}{a} \binom{s}{c}}{\binom{N}{n}}.$$

OPGAVE 14.1.1. Bewijs de formules (14.1.3) en (14.1.4).

Zijn nu r en s de bij het beschouwde experiment door $\underline{r}$ en $\underline{s}$ aangenomen waarden, dan maken wij gebruik van de voorwaardelijke verdeling van $\underline{a}$ met

deze waarden van $\underline{r}$ en $\underline{s}$ als voorwaarden. Deze formules (14.1.3) en (14.1.4) stellen ons dan in staat tot exacte berekening van de linker en rechter overschrijdskans van de gevonden waarde $\underline{a}$ van $\underline{a}$:

$$(14.1.5) \quad k_{\ell}(a) = \sum_{v=0}^a \frac{\binom{n}{v} \binom{m}{r-v}}{\binom{N}{r}} = \sum_{v=0}^a \frac{\binom{r}{v} \binom{s}{n-v}}{\binom{N}{n}},$$

resp.

$$(14.1.6) \quad k_r(a) = \sum_{v=a}^n \frac{\binom{n}{v} \binom{m}{r-v}}{\binom{N}{r}} = \sum_{v=a}^r \frac{\binom{r}{v} \binom{s}{n-v}}{\binom{N}{n}}.$$

OPMERKING 14.1.1. De grenzen, waartussen a zich kan bewegen, zijn

$$(14.1.7) \quad \max(n-s, 0) \leq a \leq \min(n, r).$$

Buiten deze grenzen worden de termen van de in (14.1.5) en (14.1.6) vermelde sommen gelijk aan 0.

OPGAVE 14.1.2. Bereken de exacte linkeroverschrijdskans, behorende bij de volgende 2×2 -tabel:

	succes	mislukking
eerste reeks	2	5
tweede reeks	3	2

(Oplossing: $0,001 + 0,044 + 0,265 = 0,310$.)

Deze exacte berekeningswijze is alleen dan zonder veel rekenwerk uitvoerbaar, als de aantallen klein zijn. Bij grotere aantallen wordt een benaderingsmethode gebruikt.

Men kan bewijzen, dat $\underline{a}$ voor m, n, r en $s \rightarrow \infty$ *asymptotisch normaal* verdeeld is als H_0 juist is.

Volgens (9.26) en (9.33) geldt als H_0 juist is

$$(14.1.8) \quad E(\underline{a} | \bar{r}=r) = \frac{nr}{N}$$

en

$$(14.1.9) \quad \sigma^2(\underline{a} | \bar{r}=r) = \frac{mnrs}{N^2(N-1)}.$$

De asymptotische normaliteit houdt in, dat de grootheid

$$(14.1.10) \quad \underline{a}^* = \frac{\underline{a} - \frac{nr}{N}}{\sqrt{\frac{mnrs}{N^2(N-1)}}}$$

bij benadering $N(0,1)$ -verdeeld is.

Bij de berekening van overschrijdingskansen past men dan bovendien nog een continuïteitscorrectie $+\frac{1}{2}$ of $-\frac{1}{2}$ toe.

VOORBEELD 14.1.1. Men vermoedt, dat bespuiting met een bepaald chemisch preparaat het voorkomen van een bepaalde ziekte onder selderieplanten zal verminderen. Om te onderzoeken of dit waar is neemt men twee groepen van 100 planten en behandelt één van deze beide groepen met het middel, de andere niet. Van de behandelde planten vertonen 20 de ziekte, van de onbehandelde 46. Betekent dit dat het middel werkt?

Oplossing: De 2×2 -tabel is:

	ziek	niet ziek	totaal
behandeld	20	80	100
onbehandeld	46	54	100
totaal	66	134	200

Onder H_0 geldt nu:

$$a = 20; \quad E(\underline{a} | \underline{r}=66) = \frac{100 \cdot 66}{200} = 33;$$

$$\sigma^2(\underline{a} | \underline{r}=66) = \frac{10^4 \cdot 66 \cdot 134}{4 \cdot 10^4 \cdot 199} = 11,11;$$

$$\sigma(\underline{a} | \underline{r}=66) = 3,33.$$

Voor de tweezijdige overschrijdingskansen berekenen wij dus

$$\frac{|\underline{a} - \frac{nr}{N}| - \frac{1}{2}}{\sqrt{\frac{mnrs}{N^2(N-1)}}} = \frac{12,5}{3,33} \approx 3,8.$$

De tweezijdige overschrijdingskans is zeer klein, zodat H_0 verworpen wordt ten gunste van de hypothese, dat de behandeling het voorkómen van de ziekte

vermindert.

OPMERKINGEN.

14.1.2. In een geval als het onderhavige kan men van mening verschillen over de vraag of men linkseenzijdig of tweezijdig moet toetsen. Gaat het er alleen om of men tot uitvoering van het middel wil besluiten of niet, dan kan men met linkseenzijdige toetsing volstaan. Indien men echter ook de mogelijkheid onder het oog ziet, dat de planten door de behandeling een verminderde weerstand tegen de ziekte gaan vertonen en stelt men, uit wetenschappelijke interesse, ook belang in die mogelijkheid, dan is tweezijdige toetsing aan te bevelen. In dit geval leiden trouwens beide methoden tot dezelfde conclusie, omdat het resultaat zeer duidelijk is.

14.1.3. De uit het experiment getrokken conclusie is alleen betrouwbaar indien de proef met de nodige voorzorg is opgezet. Is de ziekte bijv. besmettelijk en is het preparaat niet een geneesmiddel voor zieke planten, maar een afweermiddel tegen besmetting, dan dienen de 200 planten alle gelijke besmettingskansen te hebben. Dit houdt in, dat de twee groepen niet in twee bakken of velden van elkaar gescheiden geplaatst mogen zijn. Immers dan zou een toevallig iets zwaardere besmetting op het niet behandelde veld het grotere aantal zieke planten aldaar veroorzaakt kunnen hebben. In statistisch idioom: dan zijn de proeven binnen ieder der twee waarnemingsreeksen niet onafhankelijk, zodat de toets zijn geldigheid verliest. Hoe kan men dan in zo'n geval de proef zo opzetten, dat de voorwaarden wel vervuld zijn? Dit zou - met excuses voor de bewerkelijkheid van de proef - bijv. als volgt kunnen geschieden. De 200 planten worden ieder in een bakje geplant (of gezaaid), waarin nog een tweede plant geplaatst kan worden. Deze 200 bakjes, alle met dezelfde aarde gevuld, worden zo opgesteld, dat ze zoveel mogelijk aan dezelfde weersomstandigheden zijn blootgesteld. Door loting worden 100 bakken voor behandeling aangewezen en deze behandeling wordt uitgevoerd. Daarna wordt in ieder der bakken, op gelijke afstand van de proefplant, een zieke plant geplaatst. Alle zieke planten worden zoveel mogelijk van gelijk formaat en van gelijke ziektegraad genomen en zij worden ofwel vóór de bovengenoemde loting ieder aan een bepaalde bak toegewezen, ofwel door een nieuwe loting over de bakken verdeeld.

De aselechte verdeling van de zieke planten en van de behandeling over de bakken, tezamen met het voorkómen van besmetting van één bak naar een

andere (een boven nog niet vermelde, maar essentiële voorwaarde), garanderen dat een eventueel optredend systematisch verschil in ziektefrequentie ($p_1 \neq p_2$) inderdaad alleen een gevolg kan zijn van het enige systematische verschil dat nu nog tussen de beide groepen bestaat: het al of niet behandelen. De kans op een verkeerde conclusie van de vorm $p_1 < p_2$ of $p_1 > p_2$ is dan dus inderdaad gelijk aan de onbetrouwbaarheid van de toets. Worden dergelijke voorzorgen *niet* genomen, dan is deze kans echter volkomen onbekend.

Wij vermelden de volgende *eigenschappen* van deze toets zonder bewijs.

- a) De toets bezit maximaal onderscheidingsvermogen, d.w.z. men kan geen toets ontwerpen met een grotere kans op verwerping van H_0 indien $p_1 \neq p_2$ is. Het onderscheidingsvermogen is een functie van p_1 , p_2 , m en n van een vrij ingewikkeld karakter. Bij gegeven p_1 en p_2 (met $p_1 \neq p_2$) en N is het maximaal als $m = n = \frac{1}{2}N$. Bij het opzetten van proeven kan men daar rekening mee houden.
- b) De eenzijdige toetsen zijn zuiver.
- c) De toets is asymptotisch (voor m en $n \rightarrow \infty$) onderscheidend voor alle alternatieven (p_1, p_2) met $p_1 \neq p_2$.

Dezelfde toets kan ook gebruikt worden als *toets voor de onafhankelijkheid van twee kenmerken*. Beschouwen wij een populatie, waarvan de elementen de kenmerken A en B al of niet bezitten, en nemen wij daaruit een steekproef van N elementen, dan kunnen wij de resultaten rangschikken in de volgende 2×2 -tabel.

Tabel 14.1.2.

2×2 -tabel voor onafhankelijkheidstoets

	A	niet A	totaal
B	<u>a</u>	<u>b</u>	<u>m</u>
niet B	<u>c</u>	<u>d</u>	<u>n</u>
totaal	<u>r</u>	<u>s</u>	N

Hierin zijn nu alle randtotalen stochastisch, behalve N ($= \underline{m} + \underline{n} = \underline{r} + \underline{s}$).

De te toetsen hypothese is nu:

$$(14.1.11) \quad H_0: A \text{ en } B \text{ stochastisch onafhankelijk,}$$

$$\text{of: } P(A|B) = P(A|\bar{B}) = P(A).$$

Onder de voorwaarde $\underline{m} = m$ (dus ook $\underline{n} = N - m$) hebben wij nu twee reeksen waarnemingen, waarin als H_0 juist is, $P(A)$ dezelfde waarde bezit. Onder de voorwaarde $\underline{m} = m$ en $\underline{r} = r$ bezit $\underline{a}$ dus weer de hypergeometrische verdeling (14.1.3). De toets kan dus onveranderd toegepast worden.

OPGAVE 14.1.3. In een steekproef van gezinnen werd de oogkleur van man en vrouw volgens een bepaald criterium geclassificeerd als "licht" of "niet-licht". De volgende gegevens werden verkregen:

man en vrouw beiden lichte oogkleur:	309
man licht, vrouw niet-licht	: 214
man niet-licht, vrouw licht	: 133
beiden niet-licht	: 119 .

Onderzoek of de oogkleuren van man en vrouw - voor zover het de splitsing in "licht" en "niet-licht" betreft - stochastisch onafhankelijk zijn of niet.

Een derde mogelijkheid voor toepassing van de 2×2 -tabel is als *mediaan-toets*. Laat

(14.1.12) $\underline{x}_1, \dots, \underline{x}_m$ en $\underline{y}_1, \dots, \underline{y}_n$

twee onafhankelijke steekproeven zijn uit twee verdelingen, terwijl wij de hypothese H_0 willen toetsen dat deze verdelingen identiek zijn. Wij kunnen dan als volgt te werk gaan.

Laat $x_1, \dots, x_n$ en $y_1, \dots, y_n$ de waargenomen waarden voorstellen en zij M de mediaan van de gecombineerde steekproef $x_1, \dots, y_n$. Als $N = m + n$ oneven is dan is M de middelste waarneming naar grootte en als N even is het gemiddelde van de twee middelste waarnemingen.

Laat voorlopig voor het gemak N even zijn en stel dat geen van de waarnemingen gelijk is aan M . De volgende 2×2 -tabel kan nu worden opgesteld.

Tabel 14.1.3.

2×2 -tabel voor mediaantoets (N even, geen waarnemingen = M)

	aantal $> M$	aantal $< M$	totaal
eerste steekproef	<u>a</u>	<u>b</u>	m
tweede steekproef	<u>c</u>	<u>d</u>	n
totaal	$\frac{1}{2}N$	$\frac{1}{2}N$	N

Een intuïtief aannemelijke toetsingsgrootheid voor H_0 is $\underline{a}$, het aantal waarnemingen uit de eerste steekproef die groter dan M zijn. Als voor de gevonden waarde M geldt: $P(\underline{x} > M) > P(\underline{y} > M)$, dan mag worden verwacht dat $\underline{a}$ een grote waarde zal aannemen en als $P(\underline{x} > M) < P(\underline{y} > M)$, dan mag worden verwacht dat $\underline{a}$ klein zal worden.

I.p.v. $\underline{a}$ kunnen ook $\underline{b}$, $\underline{c}$ of $\underline{d}$ als toetsingsgrootheid worden gebruikt, daar i.v.m. de vaste randtotalen, de waarde van een van deze stochasten de waarden van de drie overigen volledig bepaalt.

Als H_0 waar is bezit $\underline{a}$ onder de voorwaarde $\underline{x}_1 = x_1, \dots, \underline{x}_m = x_m$, $\underline{y}_1 = y_1, \dots, \underline{y}_n = y_n$ de hypergeometrische verdeling (14.1.4), te weten:

$$P(\underline{a} = a | \underline{x}_1 = x_1, \dots, \underline{x}_m = x_m, \underline{y}_1 = y_1, \dots, \underline{y}_n = y_n) = \frac{\binom{\frac{1}{2}N}{a} \binom{\frac{1}{2}N}{b}}{\binom{N}{m}}.$$

Als namelijk H_0 waar is, de twee verdelingen zijn identiek, kan men $\underline{x}_1, \dots, \underline{x}_m, \underline{y}_1, \dots, \underline{y}_n$ beschouwen als één steekproef uit één, zij het onbekende, verdeling.

Gegeven de N trekkingen $x_1, \dots, x_m, y_1, \dots, y_n$ uit deze door H_0 gespecificeerde ene verdeling zijn er $\binom{N}{m}$ mogelijke combinaties van m elementen die als "eerste steekproef" beschouwd kunnen worden en deze hebben onder H_0 alle dezelfde voorwaardelijke kans, t.w. $\binom{N}{m}^{-1}$. Het aantal daarvan dat a elementen bevat met het kenmerk " $> M$ " en dus $b = m - a$ met het kenmerk " $< M$ " is $\binom{\frac{1}{2}N}{a} \binom{\frac{1}{2}N}{b}$.

Als N oneven is en er geen andere waarnemingen gelijk aan de mediaan M van de gecombineerde steekproef zijn, laat men M weg uit de N waarnemingen $x_1, \dots, x_m, y_1, \dots, y_n$ en gaat verder als boven te werk.

De kans op $\underline{a} = a$ onder H_0 gegeven de waarnemingen wordt dan

$$\frac{\binom{\frac{1}{2}N - \frac{1}{2}}{a} \binom{\frac{1}{2}N - \frac{1}{2}}{b}}{\binom{N-1}{m-1}} \quad \text{of} \quad \frac{\binom{\frac{1}{2}N - \frac{1}{2}}{a} \binom{\frac{1}{2}N - \frac{1}{2}}{b}}{\binom{N-1}{m}},$$

al naar gelang M wel of niet een element van de eerste steekproef is.

Als er tengevolge van diskreetheid van de verdelingen der stochasten $\underline{x}$ en $\underline{y}$ of tengevolge van onnauwkeurigheden bij het waarnemen meer waarnemingen voorkomen die gelijk zijn aan M , is er sprake van 3 kenmerken " $> M$ ", " $= M$ " en " $< M$ ". Mogelijke oplossingen in deze situatie zijn het weglaten uit de 2×2 -tabel van alle waarnemingen gelijk aan M , of een andere tweedeling der waarnemingen toepassen t.w. " $\leq M$ " en " $> M$ " of

"< M" en "> M" en wel zo, dat de twee delen in aantal zo weinig mogelijk van elkaar verschillen. Trouwens, bij elke scheiding der waarnemingen in twee klassen, mits deze scheiding slechts afhangt van de beide steekproeven tezamen en niet afzonderlijk, kan men op grond van kleine overschrijdingskansen de bovenvermelde nulhypothese H_0 van gelijkheid der verdelingen verwerpen. De afleiding van de voorwaardelijke verdeling der toetsingsgrootte $\underline{a}$ onder H_0 blijft natuurlijk hetzelfde, zodat hier

$$P(\underline{a}=\underline{a} | \underline{x}_1=\underline{x}_1, \dots, \underline{x}_m=\underline{x}_m, \underline{y}_1=\underline{y}_1, \dots, \underline{y}_n=\underline{y}_n) = \frac{\binom{A}{\underline{a}} \binom{B}{\underline{b}}}{\binom{N}{\underline{m}}},$$

waar A en B staan voor de aantallen in de twee deelklassen van de gecombineerde steekproef.

VOORBEELD 14.1.2. Bij twee groepen studenten (van studierichtingen A en B) is de studieduur als volgt verdeeld:

aantal jaren	5	6	7	8	9	10	>10	totaal
richting A	4	10	25	24	12	9	13	97
richting B	1	7	21	30	15	8	20	102
totaal	5	17	46	54	27	17	33	199

In dit geval is de mediaan 8 jaar, en deze valt middenin een groep van gelijke waarnemingen. Wij brengen daarom de scheiding naast deze groep aan en wel zodanig, dat de totalen links en rechts van deze scheiding zo weinig mogelijk van elkaar verschillen. Dit geeft ons de volgende 2×2 -tabel.

aantal jaren	≤ 8	> 8	totaal
richting A	63	34	97
richting B	59	43	102
totaal	122	77	199

Nu is onder H_0 :

$$E \underline{a} = \frac{97 \cdot 122}{199} \approx 59,5$$

en

$$\sigma^2(\underline{a}) = \frac{97 \cdot 102 \cdot 122 \cdot 77}{199^2 \cdot 198} \approx 11,85 ,$$

dus

$$\sigma(\underline{a}) \approx 3,44.$$

Dus, voor tweezijdige toetsing:

$$\frac{|63 - 59,5| - 0,5}{3,44} = \frac{3}{3,44} < 1,$$

zodat H_0 niet verworpen wordt.

OPMERKINGEN.

14.1.4. Het wordt aan de lezer overgelaten om aan te tonen dat weglating van de mediaan uit de 2×2 -tabel in bovenstaand voorbeeld ook niet tot verwerping van H_0 leidt.

14.1.5. Uit deze toepassing zien wij, dat de toets toegepast kan worden ondanks het feit, dat de precieze studieduur van hen, die meer dan 10 jaar nodig hadden, niet opgegeven is. In feite is de methode ook toepasbaar indien de gegevens niet kwantitatief, maar kwalitatief zijn: alleen van de volgorde naar grootte wordt n.l. gebruik gemaakt, niet van de grootte der waarnemingen zelf.

14.1.6. Over de vorm van de verdeling van het aantal studiejaren is geen veronderstelling gemaakt. Wat deze ook is, de toets kan steeds worden toegepast. Een dergelijke toets heet *verdelingsvrij* (beter ware: verdelingsvormveronderstellingsvrij, maar dat is wat lang) of ook (minder juiste benaming): *parametervrij*.

14.1.7. De toets is asymptotisch onderscheidend voor alle alternatieven, waarbij de medianen der twee verdelingen verschillend zijn en niet asymptotisch onderscheidend als de medianen gelijk zijn, ook al zijn de verdelingen in andere opzichten verschillend van elkaar. In dat geval is echter de verdeling van $\underline{a}$ niet meer de hypergeometrische en is de onbetrouwbaarheid van de toets onbekend. Indien de verdelingen echter in een omgeving van de (voor beide gelijke) mediaan dezelfde zijn blijft de toets

asymptotisch geldig.

Tenslotte zij nog vermeld, dat de toets vaak in iets andere vorm gegeven wordt. Als toetsingsgrootte wordt vaak opgegeven:

$$(14.1.13) \quad \chi^2 = \frac{N(|\underline{ad} - \underline{bc}| - \frac{1}{2}N)^2}{mnrs},$$

die dan, onder H_0 , een z.g. χ^2 -verdeling bezit (χ^2 = chi-kwadraat). Deze toets is, daar $(N-1)\underline{\chi}^2 = N\underline{a}^*$, behoudens een faktor $(N-1)/N$, equivalent met de boven beschreven normale benadering (met continuïteitscorrectie).

Ook wordt vaak gewerkt met het verschil der fracties successen

$$(14.1.14) \quad \frac{a}{n} - \frac{b}{m},$$

waarbij dan eveneens een normale verdeling als benadering gebruikt wordt. De verwachting, onder H_0 , van dit verschil is 0, maar de spreiding is nog afhankelijk van p_1 ($= p_2$). Voor deze spreiding wordt dan, op de in hoofdstuk XII beschreven wijze, een schatting genomen. Deze methode is daarom minder betrouwbaar dan de hier beschrevene, waarbij p_1 ($= p_2$) geëlimineerd wordt.

14.2. De toetsen van Student

W.S. GOSSET (1876-1937), als brouwer in dienst bij Guinness in Dublin, heeft - schrijvende onder de schuilnaam STUDENT - een aantal voor de moderne statistiek zeer belangrijke artikelen gepubliceerd.*¹⁾ Zijn meest bekende resultaat is de afleiding van de *verdeling van Student*, die ten grondslag ligt aan de naar hem genoemde toetsen.

Deze verdeling van Student is in feite een familie van verdelingen, die beschreven kan worden door de volgende formule voor de verdelingsdichtheid.

DEFINITIE 14.2.1. De grootte t_v heeft een Student-verdeling "met v vrijheidsgraden" als de verdelingsdichtheid is:

*¹⁾ *Student's collected papers*, edited by E.S. PEARSON & J. WISHART, University College, London, 1943.

$$(14.2.1) \quad f(t) = \frac{1}{\sqrt{v\pi}} \frac{\Gamma\left(\frac{v+1}{2}\right)}{\Gamma\left(\frac{v}{2}\right)} \left(1 + \frac{t^2}{v}\right)^{-(v+1)/2} \quad (v = 1, 2, \dots),$$

waarin Γ de (volledige) Γ -functie voorstelt.

De verdelingen van deze familie zijn alle symmetrisch t.o.v. nul en voor $v \rightarrow \infty$ teedt de $N(0,1)$ -verdeling als limiet op. Voor eindige v is de verdeling vlakker dan de $N(0,1)$ -verdeling, zoals in fig. 14.2.1 geschetst is.

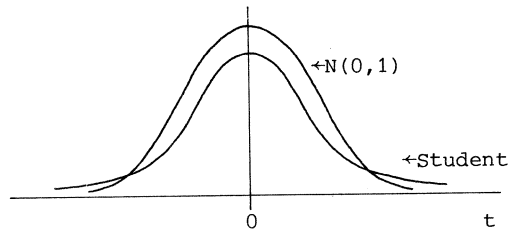


Fig. 14.2.1. Student-verdeling in vergelijking met $N(0,1)$.

In tabel 2 vindt men voor $v = 1(1)30$ de bij een aantal tweezijdige overschrijdingskansen behorende waarden van t_v . In de laatste regel staan de overeenkomstige waarden voor de $N(0,1)$ -verdeling vermeld. Met behulp daarvan kan men verifiëren, dat het in fig. 14.2.1 geschetste verschijnsel optreedt.

De toepassingen van deze verdeling berusten nu op de volgende van GOSSET afkomstige stellingen.

STELLING 14.2.1. Indien $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ onafhankelijk $N(\mu, \sigma)$ -verdeeld zijn en

$$(14.2.2) \quad \bar{\underline{x}} \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n \underline{x}_i, \quad \underline{s}^2 \stackrel{\text{def}}{=} \frac{1}{n-1} \sum_{i=1}^n (\underline{x}_i - \bar{\underline{x}})^2,$$

dan bezit de grootheid

$$(14.2.3) \quad \underline{t}_{n-1} \stackrel{\text{def}}{=} \frac{\bar{\underline{x}} - \mu}{\underline{s}} \sqrt{n}$$

een Student-verdeling met $n-1$ vrijheidsgraden.

STELLING 14.2.2. Indien $x_1, x_2, \dots, x_m$ en $y_1, y_2, \dots, y_n$ onafhankelijk $M(\mu_1, \sigma)$ resp. $N(\mu_2, \sigma)$ verdeeld zijn en

$$(14.2.4) \quad \bar{x} \stackrel{\text{def}}{=} \frac{1}{m} \sum_{i=1}^m x_i, \quad \bar{y} \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n y_i,$$

$$(14.2.5) \quad s_1^2 \stackrel{\text{def}}{=} \sum_{i=1}^m (x_i - \bar{x})^2, \quad s_2^2 \stackrel{\text{def}}{=} \sum_{i=1}^n (y_i - \bar{y})^2,$$

dan bezit de grootheid

$$(14.2.6) \quad t_{-m+n-2} \stackrel{\text{def}}{=} \frac{\bar{x} - \bar{y} - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2}{m} + \frac{s_2^2}{n}}} \sqrt{\frac{mn(m+n-2)}{m+n}}$$

een Student-verdeling met $m+n-2$ vrijheidsgraden.

OPMERKINGEN.

14.2.1. Het belang van deze stellingen is, dat de parameter σ noch in de formule voor de t 's (de indices worden vaak weggelaten), noch in de verdeling (14.2.1) optreedt. Deze parameter is dus geëlimineerd, waardoor het mogelijk wordt de aandacht volledig op μ resp. $\mu_1 - \mu_2$ te richten.

14.2.2. De veronderstelling van *normaliteit* is essentieel, evenals - bij stelling 14.2.2 - de gelijkheid van σ voor beide reeksen stochastische grootheden. Niet al te grote afwijkingen van de normaliteit, in het bijzonder indien de symmetrie der verdelingen bewaard blijft, hebben een vrij geringe invloed op de verdeling der t 's. Ongelijkheid der spreidingen van de x_i en y_j kan daarentegen de verdeling van t_{-m+n-2} sterk beïnvloeden.

TOEPASSINGEN. Het is duidelijk, dat met behulp van deze stellingen hypothesen omtrent μ resp. $\mu_1 - \mu_2$ getoetst kunnen worden. Immers indien men stelt

$$(14.2.7) \quad H_0: \mu = \mu_0,$$

waarin μ_0 een bekend getal is, dan kan men op grond van waarnemingen $x_1, x_2, \dots, x_n$ de bijbehorende waarde t_{n-1} van t_{n-1} berekenen en de daarbij behorende één- of tweezijdige overschrijdingskans in tabel 2 opzoeken. Deze toets wordt de *toets van Student voor één steekproef* genoemd. Ook de algemenere hypothese $H_0: E x_i = \mu_{0,i}$ ($i = 1, \dots, n$, met $\mu_{0,i}$ gegeven) kan

getoetst worden door $\bar{x}_i - \mu_{0,i}$ te gebruiken i.p.v. $\bar{x}_i$ en $\mu_0 = 0$.

Beschikt men over twee onafhankelijke steekproeven $x_1, x_2, \dots, x_m$ en $y_1, y_2, \dots, y_n$ dan kan men de hypothese

$$(14.2.8) \quad H_0: \mu_1 - \mu_2 = \Delta,$$

waarin Δ een bekend getal is, toetsen door t_{m+n-2} te berekenen en de bijbehorende overschrijdingskans op te zoeken. Dit is de *toets van Student voor twee steekproeven*.

VOORBEELD 14.2.1. Volgens opgave van een leverancier, die en gros vuursteentjes levert aan een afnemer, is het gemiddelde aantal steentjes per doosje gelijk aan 1000. De afnemer telt, ter controle, de inhoud van 10 aselekt gekozen doosjes en vindt de in tabel 14.2.1 in de eerste kolom vermelde aantallen.

Tabel 14.2.1.

Aantallen vuursteentjes in 10 aselekt gekozen doosjes

x	x-995	(x-995) ²
983	-12	144
1002	7	49
998	3	9
996	1	1
1002	7	49
983	-12	144
994	-1	1
991	-4	16
1005	10	100
986	-9	81
totaal	-10	594

Op grond van deze steekproef van waarnemingen kan men - in de veronderstelling, dat het aantal vuursteentjes in een doos een normale verdeling bezit - de toets van Student voor één steekproef toepassen.

Daartoe moeten $\bar{x}$ en s berekend worden. Dit gaat het gemakkelijkst door een "voorlopig gemiddelde" a te kiezen (een rond getal, dat in de buurt van $\bar{x}$ ligt, op het oog gekozen) en gebruik te maken van de volgende formules (vgl. voorbeeld 12.1):

$$(14.2.9) \quad \bar{x} = \frac{1}{n} \sum_{i=1}^n (x_i - a) + a,$$

$$(14.2.10) \quad s^2 = \frac{1}{n-1} \left\{ \sum_{i=1}^n (x_i - a)^2 - n(\bar{x} - a)^2 \right\}.$$

Wij kiezen $a = 995$; in tabel 14.2.1 zijn de waarden van $x_i - a$ en $(x_i - a)^2$ en de sommen daarvan vermeld. Wij vinden dus

$$\bar{x} = -1 + 995 = 994 \quad \text{en} \quad s^2 = \frac{1}{9}(594 - 10) = 64,89.$$

Vullen wij deze waarden in in (14.2.3), tezamen met $\mu = 1000$ omdat wij die waarde willen toetsen, dan vinden wij

$$t_9 = \frac{994 - 1000}{\sqrt{64,89}} \sqrt{10} = -2,36.$$

Voor de afnemer zijn in de regel alleen de alternatieven $\mu < 1000$ interessant, daar hij in dat geval bij de leverancier zal reclameren, doch niet als $\mu \geq 1000$ is. Hier is dus de links-éénzijdige toetsing op zijn plaats. Kiest men $\alpha_{\ell} = 0,05$ (dus $\alpha = 0,10$) dan vindt men in tabel 2 als kritieke waarde $-1,83$ (in de tabel is alleen de absolute waarde vermeld). Daar de gevonden waarde $-2,36$ kleiner is dan deze kritieke waarde, wordt de getoetste hypothese $\mu = 1000$ verworpen ten gunste van $\mu < 1000$. De conclusie luidt dus: de leverancier verpakt te weinig vuursteentjes in zijn doosjes.

OPMERKING 14.2.3. De veronderstelling van normaliteit, die essentieel is bij de toepassing van de toets van Student, wordt in de praktijk vaak stilzwijgend en zonder controle op juistheid gemaakt. Bij een gering aantal waarnemingen is deze veronderstelling ook slechts gebrekkig verifieerbaar. Wij zullen later zien, hoe hij vermeden kan worden door toepassing van een verdelingsvrije toets.

OPGAVE 14.2.1. Het opgegeven gemiddelde gewicht van zakjes suiker is 10 gram. Weging van een steekproef van 10 zakjes geeft de volgende uitkomsten: 10,04; 10,08; 10,01; 9,97; 10,02; 9,99; 10,02; 10,03; 10,00; 10,01. Toets tweezijdig, met behulp van Student's toets en met $\alpha = 0,05$, of de opgegeven gemiddelde waarde juist is.

(Antwoord: $\bar{x} = 10,017$; $s^2 = 0,00089$; $t_9 = 1,8$. De getoetste hypothese wordt niet verworpen.)

Toepassing op verschillen van paren waarnemingen

De meeste toepassingen van de toets van Student voor één steekproef hebben betrekking op een andere en iets algemenere situatie dan de boven beschrevene.

Laat $(\underline{x}_i, \underline{y}_i)$ ($i = 1, \dots, n$) n onafhankelijke paren van onafhankelijke normaal verdeelde waarnemingen zijn, met

$$(14.2.11) \quad E\underline{x}_i = \mu_i, \quad E\underline{y}_i = \nu_i,$$

dan zijn de verschillen

$$(14.2.12) \quad \underline{v}_i \stackrel{\text{def}}{=} \underline{x}_i - \underline{y}_i$$

ook normaal verdeeld, met

$$(14.2.13) \quad E\underline{v}_i = \mu_i - \nu_i.$$

Indien de $\underline{v}_i$ alle dezelfde variantie hebben (b.v. omdat alle $\underline{x}_i$ dezelfde variantie hebben, en alle $\underline{y}_i$ ook), kan nu met behulp van de toets van Student de hypothese

$$(14.2.14) \quad H_0: \mu_i - \nu_i = \Delta_{0,i},$$

met gegeven $\Delta_{0,i}$ ($i = 1, \dots, n$) getoetst worden. De grootheden $\underline{v}_i - \Delta_{0,i}$ vormen dan n.l. onder H_0 een steekproef uit een normale verdeling met verwachting 0, zodat stelling 14.2.1 toegepast kan worden.

Gewoonlijk zijn de paren $(\underline{x}_i, \underline{y}_i)$ bepalingen in duplo van de één of andere grootheid en is $\Delta_{0,i} = 0$ ($i = 1, \dots, n$) in H_0 . De toepassingen behoeven echter niet tot dat eenvoudige geval beperkt te blijven.

Een belangrijk aspect van deze toets is, dat de μ_i onderling mogen verschillen; door niveauverschillen tussen de paren wordt de toets niet verstoord.

VOORBEELD 14.2.2. De bepaling van de isolerende eigenschappen van olie kan geschieden door een glazen buisje, waarin zich twee polen bevinden, met olie te vullen en vervolgens op de polen een spanningsverschil aan te brengen, dat men laat stijgen tot de isolatie doorbroken wordt. Men kan deze bepaling van de *doorslagspanning* met dezelfde olievulling zo vaak herhalen als men wil. Wij beschouwen speciaal de eerste twee bepalingen (meting in duplo). Geven wij de doorslagspanning bij de eerste bepaling

aan met $\underline{x}$ en die bij de tweede met $\underline{y}$, dan vindt men - hetgeen niet te ver-
bazen is - in de regel verschillende uitkomsten. Het ligt echter voor de
hand te veronderstellen, dat $\underline{x}$ en $\underline{y}$ wel dezelfde kansverdeling zullen be-
zitten, ook al is deze in de regel voor verschillende oliesoorten ver-
schillend. Indien $\underline{x}$ en $\underline{y}$ niet dezelfde kansverdeling bezitten kan men
strict genomen niet van dupo-bepalingen spreken; de volgorde van bepaling
heeft dan invloed op de uitkomst. Dit is echter geenszins uitgesloten te
achten, daar een door de olie slaande vonk ionen daarin kan achterlaten.

Bij een door J.W. YODEN en J.M. CAMERON beschreven experiment ^{*)}
werden 14 verschillende oliemonsters (één van ieder van 14 soorten olie)
gebruikt; van ieder monster werden twee vullingen onderzocht, bij iedere
vulling twee bepalingen verricht. In tabel 14.2.2 staan de verschillen
 v_1 resp. v_2 van de duplobepalingen bij de eerste resp. tweede vulling ver-
meld.

Tabel 14.2.2.

Verschillen (in K.V.) van duplobepalingen
van de doorslagspanning van oliemonsters

olie monster no.	bij eerste vulling v_1	bij tweede vulling v_2	$v_1 - v_2$
1	3	-1	4
2	2	1	1
3	4	2	2
4	4	5	-1
5	1	-2	3
6	5	1	4
7	6	3	3
8	0	0	0
9	-1	7	-8
10	1	5	-4
11	7	1	6
12	-1	-1	0
13	1	0	1
14	1	8	-7

^{*)} W.J. YODEN en J.M. CAMERON, *Use of statistics to determine precision of test methods*, Symposium on Application of Statistics, Philadelphia, 1950, pp.27-34.

De veronderstelling, dat de volgorde geen invloed heeft op de uitkomsten, komt overeen met H_0 van (14.2.14), met $\Delta_{0,i} = 0$.

Daar er geen reden is aan te nemen, dat er systematische verschillen optreden tussen de eerste en de tweede vullingen met dezelfde soort olie, laten wij dit onderscheid voorlopig vallen. Wij beschikken dan over 28 onafhankelijke waarnemingen, n.l. de in de kolommen met v_1 en v_2 aangegeven verschillen.

Bepalen wij daarvan gemiddelde en variantie, dan vinden wij:

$$\bar{v} = \frac{62}{28} = 2,214,$$

$$\bar{v}^2 = 4,903,$$

$$\sum v_i^2 = 346,$$

$$s^2 = \frac{1}{n-1} (\sum v_i^2 - n\bar{v}^2) = \frac{1}{27} (346 - 28 \cdot 4,903) = 7,73,$$

$$s = \sqrt{7,73} = 2,78,$$

$$t_{27} = \frac{\bar{v} - \mu}{s/\sqrt{n}} = \frac{2,214 - \mu}{2,78/\sqrt{28}} = 4,2.$$

Deze waarde is aanzienlijk groter dan de grootste voor $v = 27$ in tabel 2 voorkomende waarde. De overschrijdingskans (tweezijdig in dit geval) is dus veel kleiner dan 0,01 en de getoetste hypothese moet met stelligheid worden verworpen. Daar $t > 0$ is luidt de conclusie, dat er een systematisch verschil tussen de beide bepalingen is in die zin, dat de tweede doorslagspanning de neiging heeft kleiner uit te vallen dan de eerste. Dit betekent niet, dat nu voor alle oliemonsters de waarden van $\Delta_i = \mu_i - v_i$ dezelfde behoeven te zijn. Zelfs behoeven zij nog niet alle $>$ te zijn, maar dit toch wel in overwegend aantal. Dit effect kan zeer wel het gevolg zijn van ionisering door de eerste doorslag.

OPGAVE 14.2.2. Indien men nu nog wil nagaan, of dit effect verschillend is voor de eerste en de tweede vulling, dan kan men de "dubbelverschillen" $v_1 - v_2$ bepalen en daarop dezelfde toets toepassen. Indien de daling van μ_i tot v_i door de eerste doorslag bij de eerste en tweede vulling dezelfde is, vormen deze verschillen opnieuw een steekproef (nu van 14 waarnemingen) uit een verdeling met verwachting 0. Voer deze berekeningen uit.

(Oplossing: $t_{13} = 0,26$. De hypothese wordt niet verworpen.)

OPMERKINGEN.

14.2.4. De boven beschreven toets is, indien $\Delta_i = \Delta$ voor alle i , asymptotisch onderscheidend als $\Delta \neq 0$ en als de Δ_i verschillend zijn nadert het onderscheidingsvermogen asymptotisch tot 1 als

$$(14.2.15) \quad \frac{1}{\sqrt{n}} \sum_{i=1}^n \Delta_i = \frac{1}{\sqrt{n}} \sum_{i=1}^n (\mu_i - \nu_i) \rightarrow \pm \infty \quad \text{voor } n \rightarrow \infty.$$

Immers voor $n \rightarrow \infty$ zal s tot σ naderen, terwijl $\bar{v}$ nadert tot $\frac{1}{n} \sum \Delta_i$, zodat $|t| \rightarrow \infty$ als aan (14.2.15) voldaan is.

14.2.5. De hier gevolgde methode van het gebruiken van beide kolommen (v_1 en v_2) tegelijk in één toets kan tot moeilijkheden leiden indien de aantallen vullingen voor verschillende oliesoorten verschillend zijn. Indien bijv. slechts bij één olie de doorslagspanning daalt, maar bij andere stijgt, kan men toch als conclusie een daling krijgen, indien men van die ene oliesoort meer proeven (vullingen) in de toetsing betreft dan van de andere samen. Wanneer men echter met alle oliesoorten evenveel proeven doet, bestaat dit bezwaar niet. Proeven als deze berusten trouwens steeds op de gedachte, dat het (eventuele) verschil tussen de beide doorslagspanningen voor alle of nagenoeg alle oliemonsters wel hetzelfde teken zal bezitten.

VOORBEELD 14.2.3. (Voorbeeld van de toets van Student voor twee steekproeven)

Van twee soorten varkensvoer, A en B, wenst men na te gaan of één van beide meer gewicht aanzet dan het andere. Men verdeelt daartoe een groep van 24 varkens aselekt in twee groepen van 12 en geeft gedurende enige tijd aan de ene groep voedsel A en aan de andere voedsel B. De gewichtstoename, die hiervan het gevolg zijn, staan in tabel 14.2.3 vermeld.

Aannemende, dat de gewichtstoename een normale verdeling bezit met gelijke spreiding voor A en B, kan de toets van Student voor twee steekproeven toegepast worden om de hypothese te toetsen, dat de verwachting der gewichtstoename voor beide voedsels gelijk is. Daartoe wordt formule (14.2.6) gebruikt, met $\mu_1 - \mu_2 = 0$. Berekening geeft: $t_{22} = 2,65$, hetgeen overeenkomt met een overschrijdingskans (tweezijdig) tussen 0,02 en 0,01. De hypothese wordt dus verworpen (als $\alpha \geq 0,02$ gekozen was) ten gunste van voedsel A (waarvoor de gemiddelde gewichtstoename in de steekproef groter is dan bij B).

Tabel 14.2.3.

Gewichtstoenames (in K.G.) van varkens gevoed
met twee typen varkensvoer

A	B
31	26
34	24
29	28
26	29
32	30
35	29
38	32
34	26
30	31
29	29
32	32
31	28

OPMERKINGEN.

14.2.6. Bij deze toets is niet alleen de normaliteit van belang, maar vooral ook de gelijkheid der beide spreidingen. In de beide steekproeven verschillen deze niet sterk, zoals grofweg te zien is door de *spreidingsbreedte* (de grootste min de kleinste waarneming) te bepalen. Deze bedraagt voor de eerste steekproef $38 - 26 = 12$ en voor de tweede $32 - 24 = 8$. Daar beide steekproeven van dezelfde omvang zijn, zijn deze twee uitkomsten vergelijkbaar (bij ongelijke omvang is dat niet zo: hoe groter de steekproef, hoe groter ook in de regel de spreidingsbreedte is). Een verschil van 30% is voor kleine steekproeven niet verontrustend. Men kan de gelijkheid der spreidingen ook rigoreus toetsen (uitgaande van de veronderstelling van normaliteit), met behulp van het quotiënt der steekproefvarianties als toetsingsgrootte; deze toets behandelen wij echter niet.

14.2.7. Bij het opzetten van een proef als deze zijn er twee belangrijke richtlijnen ter vergroting van het onderscheidingsvermogen van de toets. Ten eerste wordt dit, onder overigens gelijkblijvende omstandigheden, maximaal, indien men de beide steekproeven van gelijke omvang neemt. Ten tweede doet men er goed aan het proefmateriaal - hier de 24 varkens - zoveel mogelijk "homogeen" te maken, d.w.z. van gelijke leeftijd, geslacht en gewicht. Alle verschillen tussen de varkens verkleinen het onderscheidingsvermogen, al wordt de onbetrouwbaarheid van de toets - door het aselekt verdelen in twee groepen - er niet door beïnvloed. Zware varkens

zijn niet voor niets zwaar; zij zijn in het verleden snel gegroeid en zullen dit tijdens de proef dus waarschijnlijk ook doen. Grote verschillen in groeisnelheid vergroten echter σ , dus ook de noemer van t , waardoor deze zelf kleiner neigt te worden dan bij varkens met kleinere verschillen in groeisnelheid. Zorgvuldige selectie van proefdieren - gevolgd door aselechte verdeling in twee groepen, voor alle zekerheid - loont daarom de moeite.

14.2.8. Tabel 14.2.3 beschouwende zou men wellicht geneigd zijn de verschillen in de regels te vormen en op deze verschillen de toets van Student voor één steekproef toe te passen. Hoewel deze toets, mathematisch gezien, niet onjuist is - de onbetrouwbaarheid is immers bekend - bestaat het bezwaar, dat het onderscheidingsvermogen ervan veel kleiner is dan dat van de toets voor twee steekproeven. De gewichtstoenamen, die in één regel staan, hebben in feite niets met elkaar te maken; vermoedelijk staan de beide kolommen opgeschreven in de volgorde van weging, zodat de samenvoeging in paren op aselechte wijze tot stand is gekomen. Het is intuïtief reeds wel duidelijk, dat dit een gebrekkige, weinig onderscheidende, wijze van toetsen is.

Indien echter geen "homogene" groep varkens beschikbaar is, kan men met vrucht terugvallen op de toets voor één steekproef van verschillen. Men kan dan n.l. paren van vrijwel gelijke varkens vormen, waarbij het ene paar aanzienlijk kan verschillen van het andere. Van ieder paar wordt dan aselekt één varken aan A, één aan B toegewezen. Deze paren verschijnen dan ook - althans hun gewichtstoenamen - ieder in één regel. Nu is de situatie juist andersom: de toets voor één steekproef, toegepast op de verschillen der gewichtstoenamen, is nu de aangewezen.

OPGAVE 14.2.3. Onderstaande twee steekproeven zijn afkomstig uit twee normale verdelingen met gelijke spreiding:

0,87	0,52
0,75	1,94
0,03	0,53
0,17	1,02
0,67	0,94
1,92	0,06
0,23	0,25
2,14	2,05
2,69	1,93
1,31	1,17
	1,38
	1,43
	1,74

Toets, met behulp van de toets van Student voor twee steekproeven, de hypothese, dat de verwachting der beide verdelingen dezelfde is. (Opgave 80 van het Opgavenboekje, onderdeel a; oplossing blz. 118-119.)

14.3. Verdelingsvrije toetsen voor dezelfde of analoge hypothesen

Zoals reeds is opgemerkt kan de onderstelling van normaliteit der verdelingen vermeden worden zonder dat daarvoor een andere veronderstelling omtrent de vorm der verdelingen in de plaats gesteld behoeft te worden. Er zijn vele dergelijke "verdelingsvrije" toetsen en wij behandelen er slechts enkele van.

Toetsen voor één steekproef

Als toets voor één steekproef kan in plaats van de toets van Student voor de hypothese (14.2.7), als bekend is dat de verdeling symmetrisch is, gewerkt worden met het aantal der waarnemingen, die $< \mu_0$ zijn, als toetsingsgrootheid. Is n.l. de verdeling symmetrisch dan impliceert $\mu = \mu_0$:

$$(14.3.1) \quad P(\underline{x} < \mu_0) = P(\underline{x} > \mu_0).$$

Is $P(\underline{x} = \mu_0) = 0$ (bijv. omdat de verdeling continu is), dan zijn de beide kansen van (14.3.1) dus gelijk aan $\frac{1}{2}$ en dit kan getoetst worden met behulp van een binomiale toets - zie hoofdstuk XIII - met $p = \frac{1}{2}$.

Is $P(\underline{x} = \mu_0) > 0$, dan volgt uit de symmetrie van de verdeling, als $\mu = \mu_0$ is:

$$(14.3.2) \quad P(\underline{x} < \mu_0 | \underline{x} \neq \mu_0) = P(\underline{x} > \mu_0 | \underline{x} \neq \mu_0) = \frac{1}{2}$$

(de lezer bewijze dit). Voor de toepassing van de toets betekent dit, dat waarnemingen, die $= \mu_0$ zijn, buiten beschouwing gelaten worden en dat met behulp van de overige waarnemingen de relatie (14.3.2) wordt getoetst.

De toepassing van deze toets wordt, voor $n \leq 100$, vergemakkelijkt door tabel 3, waarin de linker kritieke waarden voor deze toets (de zgn. *tekentoets*; deze naam wordt later duidelijk) zijn vermeld. De rechter kritieke waarden kunnen gemakkelijk berekend worden, daar zij (wegens $p = \frac{1}{2}$) opgeteld bij de linker kritieke waarden het aantal waarnemingen $\neq \mu_0$ geven.

VOORBEELD 14.3.1. Wij kunnen deze toets toepassen op de gegevens van tabel 14.2.1 (voorbeeld 14.2.1). In dat geval is $\mu_0 = 1000$ en er zijn 10 waar-

nemingen > 1000 (dit kan men evengoed gebruiken als het aantal < 1000 ; het is het gemakkelijkste het kleinste aantal van deze twee te nemen, daar dit in tabel 3 direct met de daarin vermelde kritieke waarden vergeleken kan worden) is gelijk aan 3. Volgens tabel 3 leidt dit resultaat, ook bij ééNZijdige toetsing met onbetrouwbaarheidsdrempel 0,05 (dus bij gebruik van de kolom waarboven 0,10 staat) *niet* tot verwerping van de getoetste hypothese.

OPMERKING 14.3.1. Wij vinden dus bij toepassing van de toets van Student in dit geval een andere uitkomst dan bij de tekentoets. Dit is in overeenstemming met het feit, dat in geval van normale verdelingen - en verdelingen, die daarvan slechts weinig afwijken - de eerstgenoemde toets een groter onderscheidingsvermogen heeft dan de tekentoets. Hierover later meer.

Indien de verdeling, waaruit de steekproef afkomstig is, niet symmetrisch is, gelden (14.3.1 en 2) niet, ook niet onder de hypothese $\mu = \mu_0$. Deze hypothese kan dan noch met de toets van Student - wegens de afwijking van normaliteit - noch met de tekentoets onderzocht worden. In feite is er geen enkele toets, behalve voor *grote* steekproeven, die zonder verdere onderstellingen over de vorm van de scheve verdeling van toepassing is. Men kan dan echter met de tekentoets wel een andere hypothese toetsen, n.l. dat de *mediaan* een voorgeschreven waarde bezit; althans indien de verdeling continu is, want dan is de kans op een waarde onder de mediaan gelijk aan $\frac{1}{2}$, evenals de kans op een waarde daarboven. Zoals wij in hoofdstuk XII (blz. 130 e.v.) gezien hebben is, juist voor scheve verdelingen, de mediaan vaak een geschiktere parameter dan de verwachting (of het gemiddelde).

OPGAVE 14.3.1. In een steekproef van 225 waarnemingen uit een verdeling met onbekende vorm worden 150 waarnemingen aangetroffen, die kleiner dan een zeker getal A zijn, terwijl de overige groter dan A zijn.

Toets de volgende twee hypothesen ($\alpha = 0,05$):

- a) A is de mediaan van de verdeling,
- b) A is het derde quartile, d.w.z. voor A geldt:

$$P(\underline{x} < A) = \frac{3}{4}.$$

Ga verder na welke waarden voor p wel en welke niet verworpen worden als de getoetste hypothese luidt:

$$(14.3.3) \quad P(\underline{x} < A) = p.$$

De toepassing op *verschillen van paren waarnemingen*, die in het voorgaande voor de toets van Student beschreven is, geldt analoog voor de tekentoets. Zijn $\underline{x}$ en $\underline{y}$ onafhankelijk verdeeld volgens dezelfde kansverdeling en is $\underline{v} = \underline{x} - \underline{y}$, dan geldt voor de verdeling van $\underline{v}$, als $P(\underline{v}=0) = 0$ is:

$$(14.3.4) \quad P(\underline{v} < 0) = P(\underline{v} > 0) = \frac{1}{2},$$

en als $P(\underline{v}=0) > 0$ is:

$$(14.3.5) \quad P(\underline{v} < 0 | \underline{v} \neq 0) = P(\underline{v} > 0 | \underline{v} \neq 0) = \frac{1}{2}.$$

(De lezer bewijze dit.)

De hypothese, dat voor iedere i ($= 1, \dots, n$) twee grootheden (bijv. duplobepalingen) $\underline{x}_i$ en $\underline{y}_i$ dezelfde verdeling hebben, kan nu dus getoetst worden met behulp van het aantal positieve (of negatieve) tekens der verschillen $\underline{v}_i = \underline{x}_i - \underline{y}_i$. Hieraan ontleent de toets zijn naam.

De vorm der verdelingen is nu van geen enkel belang meer; ook de spreidingen der $\underline{v}_i$ behoeven niet gelijk te zijn.

VOORBEELD 14.3.2. Beschouw de gegevens van tabel 14.2.2 (doorslagspanningen van oliesoorten).

Het aantal waarnemingen $\neq 0$ in de eerste twee kolommen (v_1 en v_2) bedraagt 25 (de nullen worden buiten beschouwing gelaten). Het aantal negatieve hieronder is 5. Volgens tabel 3 is de bijbehorende tweezijdige overschrijdingskans $\leq 0,01$, zodat de getoetste hypothese - dat de duplobepalingen aequivalent zijn - verworpen moet worden. (Hier vinden wij dus wel dezelfde uitkomst als met de toets van Student.)

Onder de 12 "dubbelverschillen" $v_1 - v_2$, die $\neq 0$ zijn, zijn er 4 negatief. Volgens tabel 3 is dit niets bijzonders; de tweezijdige overschrijdingskans is $> 0,10$, daar de dáárbij behorende kritieke waarde 2 bedraagt.

OPMERKING 14.3.2. Eén der meest opvallende verschillen tussen de tekentoets en de toets van Student is wel, dat toepassing van de eerstgenoemde zoveel gemakkelijker is en vlugger gaat dan van de laatstgenoemde. Vooral bij het verkrijgen van een snel overzicht over waarnemingsmateriaal van uitgebreide en complexe aard kunnen toetsen zoals de tekentoets zeer

nuttig zijn.

OPGAVE 14.3.2. Een chemisch analyste verricht een aantal titraties in duplo. Iedere te titreren oplossing wordt daartoe, direct na de bereiding, in twee kolfjes gedaan en terwijl de eerste titratie plaats vindt staat het tweede kolfje te wachten. Zij verkrijgt de volgende uitkomsten:

Tabel 14.3.1.

Uitkomsten van 12 in duplo uitgevoerde titraties

Oplossing nr	1 ^e titratie	2 ^e titratie
1	21,24	25,83
2	16,84	17,35
3	15,52	16,12
4	25,68	28,54
5	24,04	24,58
6	19,77	27,42
7	11,92	14,73
8	28,83	27,52
9	17,38	14,91
10	11,01	19,87
11	23,43	24,38
12	17,16	20,55

Ga na (met onbetrouwbaarheidsdrempel 0,05) of het wachten invloed heeft uitgeoefend op de uitkomst van de titraties.

(Opgave 3 van het Sommenboekje; oplossing op blz. 48-49.)

Het verschil in onderscheidingsvermogen tussen de tekentoets en de toets van Student is, als aan de voor de toets van Student nodige veronderstellingen voldaan is, nogal groot, vooral bij kleine steekproeven. Dit is begrijpelijk, daar alleen het teken van ieder der waarnemingen in de toetsing betrokken wordt, maar niet de grootte. Hierdoor gaat de informatie verloren, die in de steekproef aanwezig is. Aan dit bezwaar wordt bijv. tegemoetgekomen door de symmetrietoets van WILCOXON.

De te toetsen hypothese van deze toets luidt als volgt:

(14.3.6) H_0 : ieder der onafhankelijke stochastische grootheden $x_1, \dots, x_m$ heeft een verdeling, die symmetrisch is met betrekking tot een bekende waarde a_i ($i = 1, \dots, m$).

De waarden a_i behoeven niet gelijk te zijn, de verdelingen overigens ook evenmin. De toets wordt toegepast met behulp van de verschillen $\underline{x}_i - a_i$, die onder H_0 symmetrisch verdeeld zijn om 0. Het is duidelijk, dat deze toets gebruikt kan worden om, bij symmetrische verdelingen, de hypothese $\mu = \mu_0$ te toetsen, terwijl hij anderzijds, evenals de tekentoets, op verschillen van duplo-waarnemingen toegepast kan worden. Immers het verschil $\underline{y}$ van twee onafhankelijke stochastische grootheden $\underline{x}$ en $\underline{y}$ is, indien $\underline{x}$ en $\underline{y}$ dezelfde verdeling bezitten, symmetrisch om 0 verdeeld.

Wij geven hier een korte beschrijving van deze toets. Een uitvoerige handleiding met tabellen is gegeven door A. BENARD en CONSTANCE VAN EEDEN^{*)}.

De toetsingsgrootte $\underline{T}$ is als volgt gedefinieerd. De waarnemingen ($\underline{x}_i - a_i$; $i = 1, \dots, m$; of: $\underline{x}_i - \underline{y}_i$ als het over verschillen gaat), worden, evenals bij de tekentoets, voor zoverre zij = 0 zijn, buiten beschouwing gelaten. Noem het aantal, dat dan overblijft, n . Aan ieder van deze waarnemingen wordt nu een rangnummer toegekend naar opklimmende grootte van de absolute waarden der waarnemingen. Zijn er gelijke absolute waarden, dan delen de desbetreffende waarnemingen de hun toekomende rangnummers gelijkelijk op. Deze rangnummers worden nu voorzien van de tekens der bijbehorende waarnemingen en $\underline{T}$ is de som van de op deze wijze verkregen getallen.

Men kan nu de volgende stelling bewijzen:

STELLING 14.3.1. *Indien H_0 juist is, geldt voor $\underline{T}$:*

$$(14.3.7) \quad E\underline{T} = 0,$$

$$(14.3.8) \quad \sigma^2(\underline{T}) = \frac{3n(n+1)^2 + n^3 - D}{12} = \frac{1}{6}n(n+1)(2n+1) - \frac{1}{12}(D-n),$$

waarin D de volgende betekenis heeft: verdeel alle waarnemingen, die $\neq 0$ zijn, in groepen met gelijke absolute waarden^{**)}. Laat t_j ($j = 1, \dots, k$) de omvang van de j^e knoop voorstellen, dan is

^{*)} A. BENARD en CONSTANCE VAN EEDEN, *Handleiding voor de symmetrietoets van WILCOXON*, Rapport S 208 (M 76) van de Statistische Afdeling van het Mathematisch Centrum, Amsterdam 1956. Rapport SW 4/70, een herziene versie van rapport S 176 van D. WABEKE en C. VAN EEDEN, *Handleiding voor de toets van WILCOXON*, is ook relevant en verkrijgbaar bij het Mathematisch Centrum.

^{**) "knopen" (Engels: "ties").}

$$(14.3.9) \quad D \stackrel{\text{def}}{=} \sum_{j=1}^k t_j^3. \quad *)$$

Voor $n \rightarrow \infty$ is verder $\underline{T}$, onder H_0 , asymptotisch normaal verdeeld.

Wij bewijzen deze stelling niet volledig, maar geven slechts het principe aan.

Onder H_0 is ieder der waarnemingen symmetrisch t.o.v. 0 verdeeld, zodat iedere absolute waarde > 0 , die in de steekproef gevonden wordt, gelijke kans (nl. $\frac{1}{2}$) heeft om bij een + resp. bij een - teken te behoren. Ieder der rangnummers heeft dus gelijke kans om positief of negatief in $\underline{T}$ voor te komen. De verdeling van $\underline{T}$ kan dus verkregen worden door voor ieder der rangnummers met kans $\frac{1}{2}$ een + en met kans $\frac{1}{2}$ een - te zetten en ze dan op te tellen. De lotingen voor + of - zijn bovendien onafhankelijk.

Daar deze procedure symmetrisch is met betrekking tot + en - (hij verandert niet als + en - verwisseld worden), is $\underline{T}$, onder H_0 , symmetrisch verdeeld met betrekking tot 0, waaruit (14.3.7) direct volgt.

Het bewijs van (14.3.8) is iets ingewikkelder. Zijn de waarden der t_j , gerangschikt naar opklimmende grootte van de rangnummers, waarbij zij behoren: $t_1, \dots, t_k$, en geven wij de bijbehorende rangnummers aan met $r_1, \dots, r_k$, dan geldt dus voor de eerste rangnummers, die alle gelijk aan r_1 zijn:

$$r_1 = \frac{1}{t_1}(1+2+\dots+t_1) = \frac{1}{2}(t_1+1).$$

Verder voor r_2 (de volgende t_2 rangnummers hebben deze waarde):

$$r_2 = t_1 + \frac{1}{2}(t_2+1),$$

enz. Algemeen:

$$(14.3.10) \quad r_j = \sum_{h=1}^{j-1} t_h + \frac{1}{2}(t_j+1).$$

Daar verder de verwachting van $\underline{T}$, onder H_0 , gelijk aan 0 is, geldt:

*) Zijn er geen gelijken onder de absolute waarden, dan zijn dus alle $t_j = 1$, $k = D = n$, en de laatste term van (14.3.8) valt weg. Tabellen van $\frac{1}{6}n(n+1)(2n+1)$ en $\sqrt{\frac{1}{6}n(n+1)(2n+1)}$ zijn, voor $n = 21(1)100$, te vinden in de bovengenoemde handleiding voor deze toets.

$$\begin{aligned}\sigma^2(\underline{T}) &= \sum_{j=1}^k t_j x_j^2 = \sum_{j=1}^k t_j \left\{ \sum_{h=1}^{j-1} t_h + \frac{1}{2}(t_j+1) \right\}^2 = \\ &= \sum_{j=1}^k \left\{ t_j \left(\sum_{h=1}^{j-1} t_h \right)^2 + t_j (t_j+1) \sum_{h=1}^{j-1} t_h + \frac{1}{4} t_j (t_j+1)^2 \right\}.\end{aligned}$$

Deze nogal ingewikkelde vorm laat zich als volgt reduceren:

$$\begin{aligned}\sum_{v=1}^n v^2 &= \sum_{v=1}^{t_1} v^2 + \sum_{v=1}^{t_2} (t_1+v)^2 + \sum_{v=1}^{t_3} (t_1+t_2+v)^2 + \dots = \\ &= \sum_{j=1}^k \left\{ \sum_{v=1}^{t_j} \left(\sum_{h=1}^{j-1} t_h + v \right)^2 \right\} = \\ &= \sum_{j=1}^k \left\{ t_j \left(\sum_{h=1}^{j-1} t_h \right)^2 + 2 \sum_{h=1}^{j-1} t_h \sum_{v=1}^{t_j} v + \sum_{v=1}^{t_j} v^2 \right\} = \\ &= \sum_{j=1}^k \left\{ t_j \left(\sum_{h=1}^{j-1} t_h \right)^2 + t_j (t_j+1) \sum_{h=1}^{j-1} t_h + \frac{1}{6} t_j (t_j+1) (2t_j+1) \right\},\end{aligned}$$

daar

$$\sum_{v=1}^{t_j} v^2 = \frac{1}{6} t_j (t_j+1) (2t_j+1).$$

Het linkerlid, $\sum v^2$, is gelijk aan $\frac{1}{6}n(n+1)(2n+1)$, zodat wij voor σ^2 krijgen:

$$\sigma^2(\underline{T}) = \frac{1}{6}n(n+1)(2n+1) - \frac{1}{6} \sum_{j=1}^k t_j (t_j+1) (2t_j+1) + \frac{1}{4} \sum_{j=1}^k t_j (t_j+1)^2,$$

hetgeen door een korte herleiding in (14.3.8) overgaat.

OPMERKING 14.3.3. Bij de berekening van de variantie van $\underline{T}$ zijn de t_j als niet-stochastische grootheden behandeld, terwijl zij toch eigenlijk ook stochastisch zijn. De in (14.3.8) gebruikte notatie is dan ook niet juist; de berekende variantie is die van de *voorwaardelijke* verdeling van $\underline{T}$, onder H_0 en onder de voorwaarde: $t_1 = t_1, t_2 = t_2, \dots, t_k = t_k$. Wij zouden dus, strikt genomen, de volgende notatie moeten gebruiken:

$$(14.3.11) \quad \sigma^2(\underline{T} | t_1, t_2, \dots, t_k).$$

Opmerkenswaard is overigens, dat de volgorde der t_j er in (14.3.8) niet toe doet; alleen $D = \sum t_j^3$ komt in de formule voor. Bij de berekening van D behoeft dus op deze volgorde niet gelet te worden.

De asymptotische normaliteit van $\underline{T}$ bewijzen wij hier niet. Dit bewijs is te vinden in: CONSTANCE VAN EEDEN and A. BENARD, *A general class of distribution-free tests for symmetry containing the tests of Wilcoxon and Fisher*, Proc. Kon. Ned. Akad. v. Wet. A 60 (1957) and Indagationes Mathematicae 19 (1957) 381-391, 392-400, 401-408.

De *exacte verdeling* van $\underline{T}$, onder H_0 en onder de voorwaarde van gegeven t 's (nu is de volgorde wel van belang) is, voor kleine n , gemakkelijk af te leiden. Voor $n = 3$ bijv. en zonder gelijken onder de absolute waarden, is deze verdeling te verkrijgen door op de $2^3 = 8$ mogelijke manieren + of - tekens toe te kennen aan de rangnummers 1, 2 en 3. Dit is in tabel 14.3.2 uitgevoerd.

Tabel 14.3.2.

Afleiding van de verdeling van $\underline{T}$, onder H_0 ,
voor $n = 3$, zonder gelijken

1	+	-	+	+	-	-	+	-
2	+	+	-	+	-	+	-	-
3	+	+	+	-	+	-	-	-
T	6	4	2	0	0	-2	-4	-6

Deze 8 mogelijke uitkomsten hebben alle kans $1/8$, dus heeft de waarde 0 de kans $1/4$, terwijl 6, 4, 2, -2, -4, -6, alle kans $1/8$ hebben.

Voor de berekening van deze exacte verdeling voor grotere n zijn recursieformules opgesteld, die te vinden zijn in het boven aangehaalde artikel van Constance van Eeden en A. Benard.

Met behulp van deze methode zijn, voor $n = 4(1)20$, de *exacte kritieke waarden* van de toets berekend, voor $\alpha = 0,01; 0,02; 0,05$ en $0,10$ (tweezijdig); deze zijn in tabel 4 te vinden. De overige in tabel 4 vermelde kritieke waarden zijn gevonden met behulp van de normale benadering voor

de verdeling van $\underline{T}$ onder H_0 . Zij maken de toepassing van de toets uitermate eenvoudig, althans als er geen gelijke absolute waarden zijn. Uitgebreidere tabellen van kritieke waarden en tabellen van overschrijdingskansen zijn te vinden in de reeds eerder genoemde handleiding.

OPMERKINGEN.

14.3.4. Zijn er wel gelijke absolute waarden, dan wordt hierdoor D vergroot, dus de variantie verkleind. Dit betekent, dat men, toch werkende met de variantie zonder gelijken, tot iets te grote overschrijdingskansen komt. Het verschil is echter, als er geen grote groepen gelijken zijn, klein.

Zijn er gelijken, dan geldt tabel 4 niet meer exact, maar als het er niet te veel zijn, nog wel bij goede benadering. Twijfelt men, dan kan men naast het gebruik van de tabel de normale benadering gebruiken. Bij kleine n en veel gelijken kan men de exacte verdeling bepalen langs de boven aangegeven weg, en daarmee dan de exacte overschrijdingskans van de gevonden waarde van $\underline{T}$.

14.3.5. Uit tabel 14.3.1 blijkt, dat de intervallen tussen twee op elkaar volgende waarden, die door $\underline{T}$ aangenomen kunnen worden, de lengte 2 bezitten, althans als er geen gelijken zijn. De *continuïteitscorrectie*, bij gebruik van de normale benadering, bedraagt daarom ± 1 . Zijn er wel gelijken, dan is de zaak veel ingewikkelder. Men handhaaft dan gewoonlijk toch deze continuïteitscorrectie.

VOORBEELD 14.3.1. (Vervolg)

Wij passen de toets nu toe op de gegevens van voorbeeld 14.2.1, waarop wij reeds eerder de toets van Student en de tekentoets hebben toegepast. De waarde van T kan berekend worden als in tabel 14.3.3.

Tabel 14.3.3.

Berekening van T voor de symmetrietoets van Wilcoxon
voor de gegevens van tabel 14.2.1.

verschillen $x_i - 1000$	in volgorde van absolute grootte	rangnummer met teken
-17	+ 2	+ 2
+ 2	+ 2	+ 2
- 2	- 2	- 2
- 4	- 4	- 4
+ 2	+ 5	+ 5
-17	- 6	- 6
- 6	- 9	- 7
- 9	-14	- 8
+ 5	-17	- 9½
-14	-17	- 9½
n = 10		T = 9 - 46 = -37

Volgens tabel 4, geen rekening houdende met de gelijken, geldt

$$0,025 < k_\ell < 0,05 ,$$

hetgeen dus bij linkseenzijdige toetsing met $\alpha_\ell = 0,05$ tot verwerping van de getoetste hypothese leidt.

Voor toepassing van de normale benadering berekenen wij eerst D:

$$D = 3^3 + 1 + 1 + 1 + 1 + 1 + 2^3 = 40$$

en vervolgens (met (14.3.8)):

$$\sigma^2 = \frac{30 \cdot 11^2}{12} + 1000 - 40 = 382,5; \quad \sigma = 19,56.$$

Derhalve wordt

$$\frac{T+1}{\sigma} = \frac{-36}{19,56} = -1,84.$$

hetgeen, volgens tabel 1, geeft: $k_\ell \approx 0,033$.

Dit klopt goed met de vorige uitkomst.

OPMERKING 14.3.6. Dit voorbeeld is een illustratie van het grotere onderscheidingsvermogen van deze toets, in vergelijking met de tekentoets.

Immers, de laatstgenoemde leidde niet tot verwerping van de getoetste hypothese, de eerstgenoemde wél, evenals de toets van Student (voorbeeld 14.2.1).

OPGAVE 14.3.3. Los opgave 14.2.1 (blz. 170) op met behulp van de symmetrie-toets van Wilcoxon.

(Antwoord: $T = 34$, $n = 9$; volgens tabel 4 is bij $\alpha = 0,05$ de rechter kritieke waarde 35, zodat H_0 nog net niet verworpen wordt. Met de normale benadering: $\sigma = 16,8$; $33/16,8 = 1,96$; $k = 0,05$, dus net op het randje. Zou men in dit geval een exacte beslissing willen nemen, dan kan dit alleen op grond van de exacte verdeling, rekening houdende met de gelijken.)

OPGAVE 14.3.4. Pas de symmetrietoets van Wilcoxon toe op voorbeeld 14.2.2 en opgave 14.2.2.

(Antwoord: voor de 25 verschillen $\neq 0$ van tabel 14.2.2 is $T = 251$, $D = 1411$, $\sigma = 73,5$, $T_r = 193$ voor $\alpha = 0,01$; voor de 12 dubbelverschillen is $T = 12$, $T_r = 44$ voor $\alpha = 0,10$.)

Toetsen voor twee steekproeven

Een verdelingsvrije toets voor twee steekproeven, de *mediaantoets*, is reeds in §14.1 behandeld. Een tweede, de *twee-steekproeven-toets van Wilcoxon*, zullen wij nu bespreken.

(14.3.12) Zijn $x_1, x_2, \dots, x_m$ en $y_1, y_2, \dots, y_n$ twee onafhankelijke steekproeven, dan houdt de *getoetste hypothese* H_0 in, dat al deze $m+n$ stochastische grootheden dezelfde kansverdeling bezitten. Aan de vorm van deze kansverdeling wordt geen enkele voorwaarde opgelegd.

Een uitvoerige handleiding (met tabellen) voor deze toets is:

DORALINE WABEKE en CONSTANCE VAN EEDEN, *Handleiding voor de toets van Wilcoxon*, Rapport SW 4/70 van het Mathematisch Centrum, Amsterdam, 1970, waarin ook verdere literatuur is vermeld.

Een korte beschrijving van de toets volgt hieronder.

De *toetsingsgrootheid* W is gelijk aan de som van

- a) tweemaal het aantal paren (x_i, y_j) met $x_i > y_j$, en
- b) het aantal paren (x_h, y_k) met $x_h = y_k$,
met $i, h = 1, \dots, m$; $j, k = 1, \dots, n$.*)

*) Vaak wordt ook de grootheid $U = \frac{1}{2}W$ als toetsingsgrootheid gebruikt.

De toets berust verder op de volgende stelling.

STELLING 14.3.2. *Indien H_0 juist is, geldt voor $\underline{W}$:*

$$(14.3.13) \quad \underline{EW} = mn,$$

$$(14.3.14) \quad \sigma^2(\underline{W}) = \frac{mn(N^3 - D)}{3N(N-1)}, \quad *)$$

met

$$(14.3.15) \quad N = m + n, \quad D = \sum_{j=1}^k t_j^3.$$

De t_j geven weer de omvangen der knopen aan; nu echter niet van de absolute waarden der waarnemingen, maar van de waarnemingen zelf. Voor m en $n \rightarrow \infty$ is $\underline{W}$ verder, onder H_0 , asymptotisch normaal verdeeld.

De formule voor σ^2 geldt, evenals bij de symmetrietoets van Wilcoxon, voorwaardelijk voor gegeven waarden $t_1, t_2, \dots, t_k$ (als de N waarnemingen in totaal uit k knopen bestaan).

Het bewijs van (14.3.13) is zeer eenvoudig. Is nl. H_0 vervuld, dan geldt voor ieder paar (i, j) (vgl. (14.3.5)):

$$P(\underline{x}_i > \underline{y}_j | \underline{x}_i \neq \underline{y}_j) = P(\underline{x}_i < \underline{y}_j | \underline{x}_i \neq \underline{y}_j) = \frac{1}{2},$$

zodat, als $\underline{x}_i \neq \underline{y}_j$, de verwachte bijdrage tot $\underline{W}$ gelijk is aan 1. Is echter $\underline{x}_i = \underline{y}_j$, d.w.z. nemen $\underline{x}_i$ en $\underline{y}_j$ dezelfde waarde aan, dan is de bijdrage tot $\underline{W}$ ook 1. Daar er mn paren (i, j) zijn, volgt (14.3.13) hieruit onmiddellijk.

De afleiding van (14.3.14) is ingewikkelder, maar eveneens geheel elementair. Het kan gevoerd worden met behulp van de reeds eerder afgeleide relaties (14.3.10) en (14.3.8). Daartoe herleiden wij eerst $\underline{W}$ tot

*) Zijn er geen gelijken ($D = N$), dan wordt deze formule:

$$(14.3.14a) \quad \sigma^2(\underline{W}) = \frac{1}{3}mn(N+1).$$

Tabellen van

$$\left. \begin{aligned} \underline{EU} &= \frac{1}{2}\underline{EW} = \frac{1}{2}mn, \\ \sigma(\underline{U}) &= \frac{1}{2}\sigma(\underline{W}) = \sqrt{\frac{1}{12}mn(N+1)} \end{aligned} \right\} \quad \underline{U} = \frac{1}{2}\underline{W}$$

zijn, voor $m = 1(1)100$, $n = 11(1)100$ en $m \leq n$, te vinden in:
Auxiliary table for Wilcoxon's two sample test, Rapport 132/S 86
van het Mathematisch Centrum, Amsterdam.

een som van rangnummers.

Wij nummeren alle $N = m + n$ waarnemingen naar opklimmende grootte en geven deze aan met

$$(14.3.16) \quad \begin{aligned} \text{aantal} & : t_1, t_2, \dots, t_k, \\ \text{rangnummer} & : r_1, r_2, \dots, r_k, \end{aligned}$$

waarbij dan weer geldt:

$$(14.3.10) \quad r_j = \sum_{h=1}^{j-1} t_h + \frac{1}{2}(t_j + 1).$$

Laat nu de rangnummers der x_i zijn: $R_1, R_2, \dots, R_m$, dan definiëren wij

$$(14.3.17) \quad R \stackrel{\text{def}}{=} \sum_{i=1}^m R_i.$$

De volgende relatie bestaat dan tussen W en R :

$$(14.3.18) \quad W = 2R - m(m+1).$$

Laat nl. onder $R_1, R_2, \dots, R_m$ voorkomen:

$$(14.3.19) \quad \begin{aligned} \text{rangnummer} & : r_1, r_2, \dots, r_k, \\ \text{aantal} & : v_1, v_2, \dots, v_k, \end{aligned}$$

dan geldt

$$(14.3.19) \quad \sum_j t_j = N, \quad \sum_j t_j r_j = \frac{1}{2}N(N+1), \quad \sum_j v_j = m.$$

Verder is dan

$$(14.3.20) \quad R = \sum_j v_j r_j$$

en

$$\begin{aligned} W &= \sum_j v_j (t_j - v_j) + 2 \sum_{j=2}^k \sum_{h=1}^{j-1} v_j (t_h - v_h) = \\ &= \left(\sum_j v_j t_j + 2 \sum_{j=2}^k \sum_{h=1}^{j-1} v_j t_h \right) - \left(\sum_j v_j^2 + 2 \sum_{j=2}^k \sum_{h=1}^{j-1} v_j v_h \right). \end{aligned}$$

De laatste van deze twee termen is gelijk aan $\left(\sum_j v_j \right)^2$, dus $= m^2$. Verder

volgt, door invullen van (14.3.10) in (14.3.20):

$$\begin{aligned}
 2R &= 2 \sum_j v_j \left\{ \sum_{h=1}^{j-1} t_h + \frac{1}{2}(t_j+1) \right\} = \sum_j v_j t_j + 2 \sum_{j=2}^k \sum_{h=1}^{j-1} v_j t_h + \sum_j v_j = \\
 &= W + m^2 + m = W + m(m+1).
 \end{aligned}$$

Wij kunnen dus i.p.v. met W ook met R werken, terwijl krachtens (14.3.18) geldt:

$$(14.3.21) \quad \underline{EW} = 2\underline{ER} - m(m+1)$$

en

$$(14.3.22) \quad \sigma^2(\underline{W}) = 4\sigma^2(\underline{R}).$$

Onder H_0 bezitten alle $\underline{x}_i$ en $\underline{y}_j$ dezelfde verdeling en zij vormen dus tezamen een steekproef $\underline{z}_1, \dots, \underline{z}_N$ uit deze verdeling, die als het ware door een aselechte verdeling in twee groepen van m en n waarnemingen is gesplitst. De voorwaarde $\underline{t}_1 = t_1, \underline{t}_2 = t_2, \dots, \underline{t}_k = t_k$, waardoor de bijbehorende verzameling (14.3.16) van rangnummers bepaald wordt, voegt aan de steekproef $\underline{z}_1, \underline{z}_2, \dots, \underline{z}_N$ deze verzameling rangnummers op ondubbelzinnige wijze toe. De stochastische grootheden $\underline{R}_1, \underline{R}_2, \dots, \underline{R}_m$ kunnen dus, onder H_0 en de voorwaarden $\underline{t}_1 = t_1, \underline{t}_2 = t_2, \dots, \underline{t}_k = t_k$ beschouwd worden als een aselechte steekproef (zonder teruglegging) uit deze verzameling van rangnummers.

Voor iedere $i = 1, 2, \dots, m$ geldt dan echter (vgl. (14.3.19a)):

$$\underline{ER}_i = \frac{1}{N} \sum_j t_j r_j = \frac{1}{2}(N+1)$$

en dus

$$(14.3.23) \quad \underline{ER} = m\underline{ER}_i = \frac{1}{2}m(N+1).$$

Tezamen met (14.3.21) wordt hieruit (14.3.13) gemakkelijk opnieuw verkregen. Verder is

$$(14.3.24) \quad \underline{ER}_i^2 = \frac{1}{N} \sum_j t_j r_j^2.$$

Beschouwen wij vervolgens de simultane verdeling van een willekeurig paar $(\underline{R}_j, \underline{R}_\ell)$, dan is dit dezelfde als van een steekproef van 2 zonder teruglegging uit de verzameling (14.3.16) van rangnummers. Dit paar kan dus $\binom{N}{2}$

verschillende paren van waarden met gelijke kansen aannemen (als de volgorde binnen het paar buiten beschouwing gelaten wordt), zodat geldt:

$$\begin{aligned} ER_{-i-l} &= \binom{N}{2}^{-1} \left\{ \sum_j \binom{t_j}{2} r_j^2 + \sum_{j=2}^k \sum_{h=1}^{j-1} t_j t_h r_j r_h \right\} = \\ &= \frac{1}{2} \binom{N}{2}^{-1} \left\{ \sum_j (t_j^2 - t_j) r_j^2 + 2 \sum_{j=2}^k \sum_{h=1}^{j-1} t_j t_h r_j r_h \right\}. \end{aligned}$$

Nu is, weer volgens (14.3.19a):

$$\frac{1}{4} N^2 (N+1)^2 = \left(\sum_j t_j r_j \right)^2 = \sum_j t_j^2 r_j^2 + 2 \sum_{j=2}^k \sum_{h=1}^{j-1} t_j t_h r_j r_h,$$

zodat wij, ook van (14.3.24) gebruik makende, krijgen:

$$ER_{-i-l} = \{N(N-1)\}^{-1} \left\{ \frac{1}{4} N^2 (N+1)^2 - NE_{-i}^2 \right\}.$$

Verder is dan

$$\begin{aligned} ER_{-i}^2 &= E \left(\sum_i R_{-i} \right)^2 = m ER_{-i}^2 + m(m-1) ER_{-i-l} = \\ &= m \left\{ \left(1 - \frac{m-1}{N-1} \right) ER_{-i}^2 + \frac{(m-1)N(N+1)}{4(N-1)} \right\} = \\ &= m \left\{ \frac{n}{N-1} ER_{-i}^2 + \frac{(m-1)N(N+1)}{4(N-1)} \right\}. \end{aligned}$$

Hierin kan nu $ER_{-i}^2 = \frac{1}{N} \sum_j t_j r_j^2$ rechtstreeks ontleend worden aan (14.3.8), indien daarin voor n gesubstitueerd wordt: N . Immers, volgens het bewijs van deze formule is $\sigma^2(\underline{T}) = \sum_j t_j r_j^2$, zodat

$$ER_{-i}^2 = \frac{1}{N} \left\{ \frac{1}{6} N(N+1)(2N+1) - \frac{1}{12} (D-N) \right\}.$$

Substitutie hiervan in de vorige vergelijking en herleiding geeft:

$$ER_{-i}^2 = \frac{mn(N^3-D)}{12N(N-1)} + \frac{1}{4} m^2 (N+1)^2,$$

zodat volgt

$$\sigma^2(\underline{R}) = ER_{-i}^2 - (ER_{-i})^2 = \frac{mn(N^3-D)}{12N(N-1)},$$

waarin $D = \sum_j t_j^3$ is.

Het bewijs van de asymptotische normaliteit van $\underline{W}$, onder H_0 , geven wij niet. Dit is, voor het geval zonder gelijke waarnemingen, te vinden in

H.B. MANN & D.R. WHITNEY, *On a test of whether one of two random variables is stochastically larger than the other*, Ann. Math. Stat. 18 (1947) 50-60

en in

D. VAN DANTZIG, Hoofdstuk 6 van de *Kadercursus Mathematische Statistiek*, Mathematisch Centrum, Amsterdam, 1947-1950.

Het bewijs van de asymptotische normaliteit als er wel gelijken optreden is impliciet te vinden in:

W.H. KRUSKAL, *A non-parametric test for the several sample problem*, Ann. Math. Stat. 23 (1952) 525-540

en in

D.J. STOKER, *Oor'n klas van toetsingsgrootheden vir die probleem van twee steekproewe*, dissertatie Amsterdam, 1955.

Bewijzen van de asymptotische normaliteit (met en zonder gelijken) zijn ook te vinden, en gemakkelijker toegankelijk, in de meer recente literatuur zoals bijvoorbeeld:

J. HÁJEK & Z. ŠIDÁK, *The theory of rank tests* (1967);

J. HÁJEK, *A course in nonparametric statistics* (1969);

H. WITTING & G. NÖLLE, *Angewandte mathematische Statistik* (1970)

en

J.D. GIBBONS, *Nonparametric statistical inference* (1971).

(Zie de literatuurlijst in de inleiding.)

De *exakte verdeling* van $\underline{W}$, onder H_0 , en onder voorwaarde van gegeven t 's (met hun volgorde) is ook nu weer op eenvoudige wijze af te leiden, indien m en n klein zijn. Daartoe schrijft men alle mogelijke steekproeven van m uit de $N = m + n$ door de t_j bepaalde rangnummers op en berekent voor iedere daarvan de door $\underline{W}$ aangenomen waarde. Iedere steekproef heeft de kans $\binom{N}{m}^{-1}$ en daaruit volgt direct de gezochte kansverdeling.

OPGAVE 14.3.5. Bereken de exacte verdeling van $\underline{W}$ voor het geval $m = 2$, $n = 3$ met $t_1 = 1$, $t_2 = 2$, $t_3 = 2$.

Ook voor deze verdeling zijn recursieformules afgeleid, nl. in het boven aangehaalde artikel van MANN & WHITNEY voor het geval zonder gelijke waarnemingen en in

L.J. SMID, *On the distribution of the test statistics of Kendall and Wilcoxon's two sample test when ties are present*,
Statistica Neerlandica 10 (1956) 205-214,

voor het geval met gelijken. (Zie ook de boven aangehaalde boeken van HÁJEK & ŠIDÁK, en van HÁJEK.)

In tabel 5 (a, b en c) zijn voor $\alpha = 0,02; 0,05$ en $0,10$ (tweezijdig) de kritieke waarden van W voor $m, n = 1(1)15$ en $m \leq n$ getabelleerd. Tabellen van exacte overschrijdingskansen, voor $m \leq n \leq 10$ zijn te vinden in de bovengenoemde handleiding voor deze toets. Een nomogram, waaruit men de overschrijdingskansen kan aflezen, samengesteld door P.J. RIJKOORT, is te vinden in: *Statistische tabellen en nomogrammen*, uitgegeven onder redactie van de Vereniging voor Statistiek door H.E. STENFERT KROESE N.V., Leiden. Uitgebreide tabellen vindt men ook in M. HOLLANDER & D.A. WOLFE, *Nonparametric statistical methods* (1973); zie de literatuurlijst in de inleiding.

VOORBEELD 14.3.2. Wij behandelen nu het probleem van voorbeeld 14.2.3 (blz. 174) met behulp van de twee-steekproeven-toets van Wilcoxon. Dit geschiedt gemakkelijk in de vorm van tabel 14.3.4.

Tabel 14.3.4.

Schema ter berekening van W

waar- neming	aantal	malen	bijdrage tot W	t_i	t_i^3
	bij A	bij B			
24	0	1	0	1	1
26	1	2	4	3	27
28	0	2	0	2	8
29	2	3	26	5	125
30	1	1	17	2	8
31	2	1	38	3	27
32	2	2	44	4	64
34	2	0	48	2	8
35	1	0	24	1	1
38	1	0	24	1	1
totaal	$m = 12$	$n = 12$	$W = 225$	$N = 24$	$D = 270$

Het is in dit geval iets eenvoudiger x en y in gedachten te verwisselen en dan W te berekenen. De uitkomst is dan 63 (de lezer controleer dit), hetgeen tezamen met de gevonden waarde, 225, gelijk is aan $2mn = 288$. Deze controle is altijd nuttig voor het ontdekken van fouten bij de berekening van W .

Storen wij ons niet aan de knopen en gebruiken wij tabel 5, dan vinden wij voor de tweezijdige overschrijdingskans

$$(14.3.25) \quad 0,02 < k < 0,05$$

(maar zeer dicht bij 0,02), hetgeen iets groter is dan op blz. met de toets van Student gevonden werd.

Berekenen wij de variantie met behulp van (14.3.14), dan vinden wij

$$\sigma^2 = \frac{144 (24^3 - 270)}{3 \cdot 24 \cdot 23} = 1178,6; \quad \sigma = 34,3$$

(de invloed der gelijken is gering; past men formule (14.3.14a) toe, dan krijgt men $\sigma = 34,6$).

Nu is

$$\frac{|W-mn| - 1}{\sigma} = \frac{225 - 144 - 1}{34,3} = 2,33 ,$$

dus $k \approx 0,02$, in goede overeenstemming met (14.3.25).

OPGAVE 14.3.6. Pas op de gegevens van opgave 14.2.3 de twee-steekproeven-toets van Wilcoxon toe.

(Opgave 80 van het Sommenboekje, onderdeel c; oplossing op blz. 119-120.)

OPGAVE 14.3.7. Behandel voorbeeld 14.1.2 met behulp van de toets van Wilcoxon. Beschouw daarbij alle waarnemingen > 10 als één knoop.

(Antwoord: $W = 11.265$; $mn = 9894$; $D = 320.371$; $\sigma^2 = 632.801$; $\sigma = 795,5$; $w-\mu-1/\sigma = 1,72$; $k = 0,085$.)

OPGAVE 14.3.8. Over het aantal verkeersongevallen met dodelijke afloop in twee steden A en B zijn de gegevens van tabel 14.3.5 beschikbaar.

Onderzoek of in één van beide steden systematisch meer ongevallen met dodelijke afloop optreden dan in de andere (onbetrouwbaarheidsdrempel 0,05). Ga hierbij uit van de veronderstelling, dat een eventueel verschil tussen A en B in ieder der onderzochte jaren in dezelfde richting ligt en

Tabel 14.3.5.

Aantal ongevallen met dodelijke afloop in
twee steden A en B

Jaar	Aantal ongevallen	
	in A	in B
1920	99	75
1921	87	83
1922	105	81
1923	97	101
1924	109	93
1925	105	97
1926	116	87
1927	110	113
1928	125	103
1929	117	93
1930	130	105
1931	137	107
1932	121	124
1933	126	112
1934	131	121
1935	145	130
1936	142	121

licht de keuze van de door U gekozen toets toe aan de hand van een grafiek van de gegevens. Wat is Uw conclusie? Welke verklaringen van het door U gevonden verschijnsel acht U aannemelijk?

(Opgave 13 uit het Sommenboekje; oplossing (gedeeltelijk) op blz. 56-58.)

OPMERKINGEN.

14.3.7. Bestaan de waarnemingen uit twee knopen, dan gaat de twee-steekproeven-toets van Wilcoxon over in de 2×2 -tabel toets van §14.1. Men vergelijkt:

CONSTANCE VAN EEDEN, *Verdelingsvrije toetsen voor twee steekproeven en de methode der 2×2 -tabel*, Statistica Neerlandica 10 (1956) 157-162.

14.3.8. Zoals in opgave 14.3.7 blijkt, is de toets van Wilcoxon voor twee steekproeven geschikt voor toepassing op afgeknot waarnemingsmateriaal. Dit is bijv. van belang bij levensduurproeven, waarbij het grootste deel der waarnemingen in betrekkelijk korte tijd verkregen wordt, terwijl men op enkele waarnemingen zeer lang moet wachten. Men kan dan beter de proef

enigszins uitbreiden (m en n groter) en na een van tevoren vastgestelde tijd afbreken (waarom "van tevoren vastgesteld"?).

OPGAVE 14.3.9.

- a) Twee soorten gloeilampen moeten met elkaar vergeleken worden wat hun levensduur betreft. Daartoe worden van de eerste soort 10 exemplaren en van de tweede 12 exemplaren genomen en deze worden gebrand tot ze doorbranden. De uitkomsten zijn (in uren):

x: 625, 637, 710, 770, 820, 843, 856, 920, 1070, 1225,

y: 630, 683, 780, 830, 889, 970, 1028, 1150, 1210, 1470, 1520, 2090.

Onderzoek of één van beide soorten een systematisch langere levensduur heeft dan de andere soort (onbetrouwbaarheidsdrempel 0,05).

- b) Indien men de boven beschreven proef wil bekorten tot maximaal 1000 uren, nemen de gegevens de volgende vorm aan:

x: 625, 637, 710, 770, 820, 843, 856, 920, 2 waarnemingen > 1000,

y: 630, 683, 780, 830, 889, 970, 6 waarnemingen > 1000.

Toets ook voor deze gegevens de hypothese, die U in deel a) van deze opgave getoetst hebt.

(Aanwijzing: beschouw de waarnemingen die groter dan 1000 zijn alle als gelijk uitgevallen waarnemingen.)

(Opgave 24 uit het Sommenboekje; oplossing op blz. 68-70.)

14.4. Onderscheidingsvermogen van de toetsen van Student, de tekentoets en de toetsen van Wilcoxon en asymptotische eigenschappen.

De berekening van het onderscheidingsvermogen hebben wij in hoofdstuk XIII voor binomiale toetsen uitgevoerd. Voor de toetsen van Student kan een soortgelijke berekening uitgevoerd worden, die echter niet tot een integraal voert, die met de normale verdeling te benaderen is, maar tot de zgn. niet-centrale t-verdeling, die evenals de normale verdeling (maar pas veel later) uitvoerig getabelleerd is. Het onderscheidingsvermogen heeft dezelfde soort vorm als bij de tekentoets, indien men op de horizontale as uitzet

$$(14.4.1) \quad \delta_1 = \frac{\mu - \mu_0}{\sigma}$$

voor de toets voor één steekproef (μ is de werkelijke verwachting, μ_0 de getoetste en σ de spreiding van $\underline{x}$) en

$$(14.4.2) \quad \delta_2 = \frac{\mu_1 - \mu_2 - \Delta}{\sigma}$$

voor de toets voor twee steekproeven ($\mu_1 - \mu_2$ is het werkelijke verschil tussen de verwachtingen van $\underline{x}$ en $\underline{y}$ en Δ de getoetste waarde daarvan; σ de gemeenschappelijke spreiding van $\underline{x}$ en $\underline{y}$).

Het onderscheidingsvermogen hangt dan af van het aantal vrijheidsgraden $v = n - 1$ bij één steekproef en $v = m + n - 2$ bij twee steekproeven.

In figuur 14.4.1 is het onderscheidingsvermogen als functie van μ geschetst in de vorm van enkele curven voor verschillende waarden van v , behorende bij $\alpha = 0,05$.

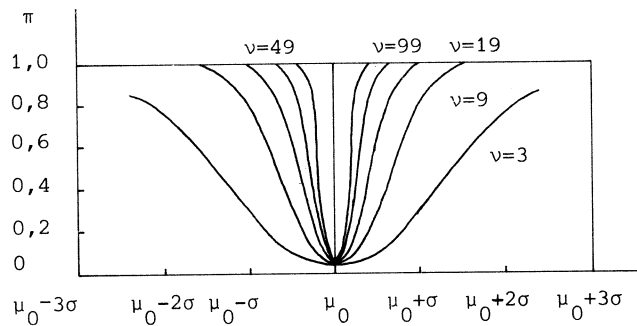


Fig. 14.4.1. Het onderscheidingsvermogen van Student's tweezijdige toetsen met $\alpha = 0,05$.

Men kan bewijzen dat de toetsen van Student, *indien de bijbehorende voorwaarden van normaliteit en gelijke spreidingen vervuld zijn*, het grootst mogelijke onderscheidingsvermogen bezitten. Dit houdt dus in, dat zij in dat geval meer onderscheidend zijn dan verdelingsvrije toetsen. Het omgekeerde kan voorkomen als aan de genoemde veronderstellingen niet is voldaan, terwijl dan bovendien de onbetrouwbaarheidsdrempel van de toetsen van Student niet meer de opgegeven waarde bezit.

De minst onderscheidende toets (in geval van normaliteit en verschuivingen als alternatieven) voor één steekproef die wij hebben behandeld, is de tekentoets. Ter illustratie zullen wij, voor $n = 20$, deze toets vergelijken met die van Student.

Wij stellen daarbij, voor de eenvoud, $\sigma = 1$ en $\mu_0 = 0$. Het onder-

scheidingsvermogen van de toets van Student kan dan uit fig. 14.4.1 afgelezen worden. Voor de tekentoets kan het als volgt berekend worden.

Voor $n = 20$ zijn de kritieke waarden bij $\alpha = 0,05$ volgens tabel 3, gelijk aan 5 en 15. Wij berekenen nu eerst, bijv. voor $\mu = 0,5$, de kans $p_{0,5}$ op een positieve uitkomst. De beschouwde verdeling is dus $N(0,5;1)$, zodat tabel 1 ons geeft: $p_{0,5} = 1 - 0,31 = 0,69$. Vervolgens berekenen wij, met behulp van de normale benadering voor de binomiale verdeling, de kans op 15 of meer positieve uitkomsten onder de 20 waarnemingen. Noemen wij het aantal positieve uitkomsten $\underline{a}$, dan is

$$E\underline{a} = 20 \cdot 0,69 = 13,8; \quad \sigma^2(\underline{a}) = 20 \cdot 0,69 \cdot 0,31 = 4,28;$$

$$\sigma = 2,07; \quad \frac{14,5 - 13,8}{2,07} = 0,34 ,$$

zodat de gezochte kans, volgens tabel 1, $\approx 0,367$ is. De kans om in het andere deel van de kritieke zone ($a \leq 5$) terecht te komen is gering en wordt buiten beschouwing gelaten. Uit fig. 14.4.1 volgt, dat de overeenkomstige kans, voor de toets van Student, $\approx 0,55$ is.

Analoge berekening voor $\mu = 1, 1,5$ en 2 leiden tot de in fig. 14.4.2 geschetste curve. In deze figuur is, ter vergelijking, de desbetreffende curve voor de toets van Student eveneens opgenomen. Voor $\mu = 0$ is $p_0 = \frac{1}{2}$ en de onbetrouwbaarheid blijkt dan, bij opzoeken in een tabel van de binomiale verdeling, 0,0414 te bedragen, dus vrij aanzienlijk minder dan die van de toets van Student, waarvoor deze exact 0,05 is. Een gedeelte van het geringere onderscheidingsvermogen is aan deze kleinere onbetrouwbaarheid te wijten.

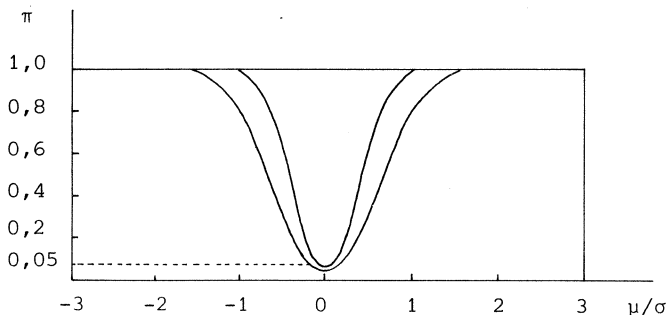


Fig. 14.4.2. Onderscheidingsvermogen van de tekentoets in vergelijking met de toets van Student (één steekproef, $\alpha = 0,05$ tweezijdig, $n = 20$).

OPGAVE 14.4.1. Voer de berekening voor fig. 14.4.2 uit.

Het onderscheidingsvermogen van de toetsen van Wilcoxon, evenals van vele andere verdelingsvrije toetsen, is minder gemakkelijk te bestuderen. De belangrijkste reden hiervan is, dat dit onderscheidingsvermogen niet verdelingsvrij is. Ook bij verschuiving der verdelingen ten opzichte van de getoetste ligging, met behoud van de vorm, is het onderscheidingsvermogen van deze vorm afhankelijk.

Verder blijkt, dat, in tegenstelling met het eenvoudige karakter van de verdeling van de toetsingsgrootheden $\underline{T}$ en $\underline{W}$ onder H_0 , deze grootheden onder alternatieve hypothesen juist zeer ingewikkelde verdelingen bezitten. Het onderzoek naar de eigenschappen van verdelingsvrije toetsen vindt dan ook op vele andere manieren plaats, zoals:

- het vergelijken van toetsen d.m.v. simulatie van steekproeftrekkingen uit bepaalde verdelingen m.b.v. de computer. Zie bijvoorbeeld

J. HEMELRIJK, *Experimental comparison of Student's and Wilcoxon's two-sample tests*, Symposium Quantitative Methods in Pharmacology (1960) North-Holland;

- de bestudering van het asymptotisch onderscheidingsvermogen. Hierbij gaat het zowel om de asymptotische verdeling van de toetsingsgrootheden onder alternatieve hypothesen, als om het onderzoek of een bepaalde toets al of niet asymptotisch meest onderscheidend is. Zie i.h.b.

J. HÁJEK & Z. ŠIDÁK, *Theory of rank tests* (1967)

en

H. WITTING & G. NÖLLE, *Angewandte mathematische Statistik* (1976);

- de bestudering van het lokale onderscheidingsvermogen. D.w.z. van het onderscheidingsvermogen in een (kleine) omgeving van een bepaalde relevante parameterwaarde, i.h.b. een omgeving van de nulhypothese, bij een gegeven funktionele vorm van de alternatieve verdelingen. Een uitvoeribehandeling van het lokale onderscheidingsvermogen van vele verdelingsvrije toetsen is te vinden in de bovengenoemde werken van HÁJEK & ŠIDÁK en WITTING & NÖLLE;
- de bestudering van de (asymptotische) relatieve efficiency. Als men twee toetsen vergelijkt m.b.t. hun onderscheidingsvermogen, kan men ook kijken naar het quotiënt n/m , waar n staat voor een willekeurig vast aantal

waarnemingen en m voor het aantal waarnemingen dat men moet verrichten om voor een bepaalde vaste alternatieve parameterwaarde met de ene toets een even groot onderscheidingsvermogen te bereiken als bij de andere toets met n waarnemingen. Dit quotiënt n/m geeft dan een indruk van de doeltreffendheid (efficiency) van de ene toets t.o.v. de andere. Deze eenvoudige vergelijkingsmethode is in verschillende richtingen ontwikkeld.

I.h.b. is de asymptotische relatieve efficiency van belang. Hierbij is men o.a. geïnteresseerd in de limiet van quotiënten n/m bij toename van deze steekproefgrootten onder bepaalde bijkomende voorwaarden m.b.t. het onderscheidingsvermogen en de alternatieve hypothesen. Veel gebruikte definities van asymptotische efficiency zijn die van PITMAN, BAHADUR en van HODGES & LEHMANN.

Men raadplege weer de bovenvermelde werken van HÁJEK & ŠIDÁK en van WITTING & NÖLLE, en de in 1977 te verschijnen MC-syllabus over efficiency-begrippen in de statistiek.

De resultaten van al deze onderzoeken kort samenvattend kan men m.b.t. het onderscheidingsvermogen van de teken- en Wilcoxon-toetsen het volgende opmerken.

In die situaties, waarin aan alle voorwaarden m.b.t. de beste parametrische toetsen (de Student-toetsen) is voldaan, hebben deze verdelingsvrije toetsen (uiteraard) een geringer onderscheidingsvermogen.

Bij het steekproeftrekken uit normale verdelingen hebben zowel de een- als de tweesteekproeven tekentoets Pitman-efficiency $\frac{2}{\pi} = 0,64$ t.o.v. de betreffende Student-toetsen. De beide toetsen van Wilcoxon hebben Pitman-efficiency $\frac{3}{\pi} = 0,95$ t.o.v. de Student-toetsen. De Wilcoxon-toetsen hebben Pitman-efficiency $\frac{3}{2}$ t.o.v. de tekentoetsen in dezelfde situatie.

Als de steekproeven daarentegen getrokken worden uit logistische verdelingen (d.w.z. verdelingen met dichtheden

$$f(x; \mu, \sigma) = \sigma^{-1} e^{-(x-\mu)/\sigma} / (1 + e^{-(x-\mu)/\sigma})^2, \quad -\infty < \mu < \infty, \quad \sigma > 0,$$

deze hebben "dikkere" staarten dan de normale verdelingen) dan zijn de Wilcoxon-toetsen lokaal en asymptotisch meest onderscheidend.

Als de steekproeven getrokken worden uit dubbel exponentiële verdelingen (d.w.z. verdelingen met dichtheden

$$f(x; \mu, \sigma) = \frac{1}{2}\sigma^{-1} e^{-|x-\mu|/\sigma}, \quad -\infty < \mu < \infty, \quad \sigma > 0,$$

deze hebben weer "dikkere" staarten dan de logistische verdelingen) dan zijn de tekentoetsen lokaal en asymptotisch meest onderscheidend.

Daar i.h.a. echter niet veel konkrete informatie beschikbaar zal zijn over de vorm van de onderliggende verdeling, zullen de Wilcoxon-toetsen in vele gevallen aan te bevelen zijn.

In het normale geval zijn zij praktisch gelijkwaardig met de Student-toetsen en in andere gevallen ("scheve" basisverdelingen of verdelingen met "dikke" staarten) vaak superieur of gelijkwaardig. De tekentoetsen zijn in de meeste gevallen inferieur aan de Wilcoxon-toetsen, maar vanwege de eenvoud van uitvoering zeer geschikt voor een eerste oriëntatie.

Er bestaan echter andere, zij het iets meer gecompliceerde verdelingsvrije toetsen, die in het normale geval asymptotisch gelijkwaardig zijn aan de Student-toetsen. Voor informatie m.b.t. deze laatstgenoemde toetsen kan de lezer de in de inleiding genoemde literatuur over verdelingsvrije methoden raadplegen en ook de reeds eerder genoemde werken van HAJEK & ŠIDÁK en WITTING & NÖLLE.

Het boek van

J.V. BRADLEY, *Distribution-free statistical tests* (1968)

geeft voor vele verdelingsvrije toetsen uitvoerige informatie over (experimenteel) onderzoek naar het onderscheidingsvermogen bij kleine steekproeven en informatie over de Pitman-efficiency van deze toetsen t.o.v. elkaar.

Van alle besproken toetsen kan, grofweg, gezegd worden, dat zij asymptotisch onderscheidend zijn, zodra de getoetste hypothese niet is vervuld. Zo is dit bijv. het geval voor de tweezijdige toetsen van Student als $\mu \neq \mu_0$ resp. $\mu_1 - \mu_2 \neq \Delta$ is, maar niet als $\mu = \mu_0$ resp. $\mu_1 - \mu_2 = \Delta$ is, terwijl de verdeling niet normaal is of de spreidingen niet gelijk zijn. Het onderscheidingsvermogen stijgt dan soms wel boven de onbetrouwbaarheidsdrempel (zoals reeds eerder signaleerd), maar het hoeft niet tot 1 te naderen.

Bij de toets van Wilcoxon voor twee steekproeven is de grootheid, waar het om gaat:

$$(14.4.3) \quad P(\underline{x} > \underline{y}) - P(\underline{x} < \underline{y}).$$

Onder H_0 is dit verschil = 0 en als het $\neq 0$ is, is de tweezijdige toets asymptotisch onderscheidend. Voor de symmetrietoets is de zaak iets ingewikkelder. Men zie voor dat geval de eerder geciteerde literatuur.

HOOFDSTUK XV

BETROUWBAARHEIDSGRENZEN

De theorie der betrouwbaarheidsgrenzen vormt de verbinding tussen de schattings- en de toetsingstheorie. Ook hier wordt het onderwerp van onderzoek gevormd door één (of meer) onbekende parameter(s) van een kansverdeling, waarover door een reeks, gewoonlijk onafhankelijke, waarnemingen

$$(15.1) \quad \underline{w}_1, \underline{w}_2, \dots, \underline{w}_n$$

informatie wordt verschaft.

DEFINITIE 15.1. Een functie $\underline{t}_\ell = t(\underline{w}_1, \underline{w}_2, \dots, \underline{w}_n)$ van de waarnemingen is een *betrouwbaarheidsondergrens* voor een onbekende parameter θ , met onbetrouwbaarheidsdrempel α_ℓ , indien geldt:

$$(15.2) \quad P(\underline{t}_\ell < \theta) \geq 1 - \alpha_\ell,$$

waarin θ de (onbekende) waarde der parameter voorstelt.

Analoog geldt, voor een *betrouwbaarheidsbovengrens* met onbetrouwbaarheidsdrempel α_r :

$$(15.3) \quad P(\underline{t}_r > \theta) \geq 1 - \alpha_r.$$

Tezamen vormen $\underline{t}_\ell$ en $\underline{t}_r$ een *betrouwbaarheidsinterval* voor θ , met onbetrouwbaarheidsdrempel α , indien geldt:

$$(15.4) \quad P(\underline{t}_\ell < \theta < \underline{t}_r) \geq 1 - \alpha.$$

OPMERKINGEN.

15.1. In alle praktisch voorkomende gevallen, waarbij voor één parameter zowel onder- als bovengrenzen zijn afgeleid, geldt:

$$(15.5) \quad P(t_{-l} < t_{-r}) = 1.$$

Indien hieraan voldaan is, geldt

$$(15.6) \quad \alpha = \alpha_r + \alpha_l.$$

(De lezer bewijze dit.)

15.2. Evenals bij toetsen kan de toevoeging "drempel" in "onbetrouwbaarheidsdrempel" vervallen, indien t_{-l} en t_{-r} een continue kansverdeling bezitten; de ongelijkheden (15.2,3 en 4) gaan dan in gelijkheden over - niet binnen, maar buiten de haken -. Zijn t_{-l} en t_{-r} discreet verdeeld, dan is de onbetrouwbaarheid hoogstens gelijk aan en in de regel kleiner dan de onbetrouwbaarheidsdrempel; de eerstgenoemde hangt van de onbekende waarde van θ af, zodat hij zelf ook onbekend blijft. Voor sommige (mogelijke) waarden van θ worden onbetrouwbaarheid en onbetrouwbaarheidsdrempel gelijk, zodat de ongelijkheden (15.2,3 en 4) niet "verbeterd" kunnen worden.

15.3. $1 - \alpha_l$, $1 - \alpha_r$ resp. $1 - \alpha$ worden *betrouwbaarheidscoefficienten* *) genoemd.

15.4. Aan het einde van hoofdstuk XII hebben wij reeds enkele voorbeelden van (benaderde) betrouwbaarheidsgrenzen ontmoet. Wij zullen in dit hoofdstuk de exacte grenzen, behorende bij die voorbeelden, afleiden.

Afleiding met behulp van de toetsingstheorie

Wij bespreken eerst het verband van de theorie der betrouwbaarheidsgrenzen met de toetsingstheorie, daarna met de schattingstheorie.

Indien wij beschikken over een toetsingsmethode T, waarmee wij, op grond van de waarnemingen (15.1), iedere toegelaten (mogelijk geachte) waarde van de onbekende parameter kunnen toetsen, dan geldt de volgende stelling.

STELLING 15.1. De verzameling der voor θ toegelaten waarden (de parameter-ruimte Ω) wordt, door toepassing van de toetsingsmethode met onbetrouwbaarheidsdrempel α , in twee deelverzamelingen gesplitst. Dit zijn:

$$(15.7) \quad \underline{G} = G(\underline{w}_1, \underline{w}_2, \dots, \underline{w}_n)$$

de verzameling van alle niet verworpen waarden, en

*) Engels: confidence coefficient, of: confidence level.

$$(15.8) \quad \bar{G} = \bar{G}(w_1, w_2, \dots, w_n) = \Omega - G$$

de verzameling van alle wel verworpen waarden, het complement van G met betrekking tot Ω . Voor G geldt:

$$(15.9) \quad P(\theta \in G) \geq 1 - \alpha,$$

waarin θ de werkelijke waarde van de onbekende parameter voorstelt (en " $\in$ " de relatie: "is een element van"). G is dus een betrouwbaarheidsgebied voor θ , met onbetrouwbaarheidsdrempel α .

BEWIJS. Indien de werkelijke waarde θ getoetst wordt, is, per definitie van α , de kans op verwerping $\leq \alpha$. Hieruit volgt (15.9) direct. $\square$

OPMERKING 15.5. Bij een toets T behoort dus een betrouwbaarheidsgebied G . Omgekeerd wordt T weer uit G verkregen door een getoetste waarde θ_0 te verwerpen, indien θ_0 niet in G ligt en niet, indien $\theta_0 \in G$.

VOORBEELD 15.1. Betrouwbaarheidsinterval voor een onbekende kans.

Wij leiden nu, volgens het in stelling 15.1 beschreven principe, een methode af om exacte betrouwbaarheidsgrenzen voor een onbekende kans te vinden.

Daarvoor is de volgende hulpstelling nodig.

LEMMA 15.1. *De rechteroverschrijdingskans*

$$(15.10) \quad P(\underline{x} \geq x; p) = \sum_{i=x}^n \binom{n}{i} p^i q^{n-i}$$

van een binomiale verdeling, behorende bij een gegeven waarde van x , is een monotoon stijgende functie van p . (De linkeroverschrijdingskans is dus een monotoon dalende functie van p .)

BEWIJS. Beschouw één term van (15.10):

$$(15.11) \quad \phi(p) \stackrel{\text{def}}{=} \binom{n}{i} p^i q^{n-i}.$$

Voor de afgeleide hiervan geldt

$$(15.12) \quad \phi'(p) = \binom{n}{i} p^{i-1} q^{n-i-1} (i - np),$$

een relatie, die voor $i = 0$ en $i = n$ zijn geldigheid behoudt. Uit (15.12) volgt

$$(15.13) \quad \phi'(p) = \begin{cases} > 0 & \text{voor } i > np, \\ = 0 & \text{voor } i = np, \\ < 0 & \text{voor } i < np. \end{cases}$$

Is nu in (15.10) $x > np$, dan bezitten alle termen een positieve afgeleide, zodat dit voor de som ook geldt. Is $x \leq np$, dan is $x-1 < np$ en dan bezitten alle termen van de complementaire som (de linkeroverschrijdingskans van x) negatieve afgeleiden, dus hun som ook. Trekken wij deze som van 1 af, dan wordt weer (15.10) verkregen, die dus ook in dit geval een positieve afgeleide bezit. $\square$

Wij beschouwen nu eerst de rechtseenzijdige toets, die, bij gegeven n en α_r , berust op een kritieke waarde x_r , die echter nog afhangt van de getoetste waarde p_0 . Deze waarde x_r is de kleinste, waarvoor geldt

$$(15.14) \quad P(\underline{x} \geq x_r; p_0) \leq \alpha_r,$$

zodat

$$(15.15) \quad P(\underline{x} \geq x_r - 1; p_0) > \alpha_r$$

is. Vergroten wij nu p_0 , dan wordt volgens lemma 15.1 het linkerlid van (15.15) bij vasthouden van x_r , eveneens vergroot, zodat hij $> \alpha_r$ blijft. Dit betekent dat x_r door vergroting van p_0 niet kan dalen, maar - uiteraard - wel kan stijgen.

Bij een gegeven waargenomen waarde x van $\underline{x}$ (het aantal "successen" in n onafhankelijke proeven; dit is in dit geval het waarnemingsmateriaal), worden nu die waarden p_0 verworpen, waarvoor $x > x_r$ is. De waarde $p_0 = 1$ wordt zeker niet verworpen, want daarvoor is de rechter kritieke zone leeg. Laten wij p_0 nu dalen vanaf de waarde 1, dan wordt voor een zekere waarde van p_0 : $x_r = n$, en x_r daalt vervolgens trapsgewijze.

Zolang echter $x_r > x$ is, wordt p_0 nog niet verworpen. Bij een zekere waarde van p_0 gaat x_r echter van de waarde $x+1$ over in de waarde x . Deze waarde van p_0 geven wij aan met p_ℓ ; het is de eerst die verworpen wordt. Alle waarden $p_0 > p_\ell$ worden niet, alle waarden $\leq p_\ell$ wel verworpen.

Volgens stelling 15.1 is dus p_ℓ , die een functie van het aantal suc-

cessen x is, de betrouwbaarheidsondergrens voor p , zodat, als wij x weer als stochastisch beschouwen, geldt:

$$(15.16) \quad P(p_\ell < p) \geq 1 - \alpha_r,$$

waarin p de (onbekende) werkelijke waarde van de kans op succes voorstelt. De onbetrouwbaarheid α_r van de rechtseenzijdige toets is dus gelijk aan die van de betrouwbaarheidsondergrens. In deze laatste betekenis zou hij juist met α_ℓ aangegeven worden.

Uit het bovenstaande blijkt:

STELLING 15.2. Een betrouwbaarheidsondergrens p_ℓ voor een onbekende kans wordt verkregen door de grootste waarde van p_0 op te zoeken, waarvoor geldt:

$$(15.17) \quad x = x_r,$$

met x = aantal successen en x_r = de bij de voor deze ondergrens gekozen onbetrouwbaarheidsdrempel behorende rechter kritieke waarde der binomiale toets.

Op analoge wijze wordt een betrouwbaarheidsbovengrens p_r voor p gevonden als de kleinste waarde van p_0 , waarvoor

$$(15.18) \quad x = x_\ell,$$

met x_ℓ = de linker kritieke waarde van de binomiale toets met de voor de bovengrens gekozen betrouwbaarheidsdrempel.

Beschikt men over exacte tabellen van de binomiale verdeling of van x_r , voor vele waarden van n en p , dan kan men dus bij gegeven n en x , de bijbehorende waarden van p_ℓ en p_r door opzoeken bepalen.

Voor $n = 10$ en $x = 0, 1, \dots, 10$ is het resultaat daarvan voor $\alpha_\ell = \alpha_r = 0,025$ in fig. 15.1 in beeld gebracht. De dik getrokken stukken der verticale lijnen stellen de betrouwbaarheidsintervallen voor, de gestippelde gedeelten de verworpen hypothesen. De verbindingslijnen der grenspunten hebben slechts aanschouwelijke betekenis. De horizontale dienen voor het gemak van de aflezing^{*)}.

^{*)} Deze figuur vindt men ook in het leerboek van Dr. L.N.H. BUNT, *Statistiek voor het voorbereidend hoger en middelbaar onderwijs*, Wolters, Groningen 1956, blz. 150. Daar wordt de hier kort besproken theorie uitvoerig en elementair uiteengezet.

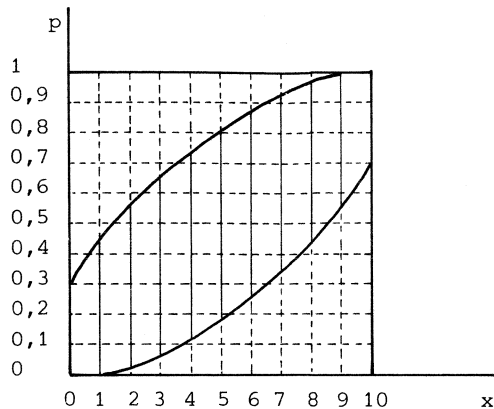


Fig. 15.1. Betrouwbaarheidsintervallen voor p bij $n = 10$ en $\alpha_l = \alpha_r = 0,025$.

Voor grotere n kan men gebruik maken van de normale benadering met continuïteitscorrectie. Dit leidt tot expliciete formules voor p_l en p_r , als functie van n , x , α_l en α_r , die echter een benaderend karakter bezitten. Zij zijn echter veel nauwkeuriger dan de benaderingsformule (12.77), die op slordige wijze is afgeleid. De beide formules zijn wel asymptotisch, voor $n \rightarrow \infty$, equivalent.

Voor de afleiding van de benaderingsformules maken wij gebruik van formules (13.42) en (13.43) voor x_l en x_r , die op de normale benadering met continuïteitscorrectie berusten. De tweede hiervan luidt:

$$(15.19) \quad x_r \approx np_0 + u_{\alpha_r} \sqrt{np_0(1-p_0)} + \frac{1}{2},$$

waarin u_{α_r} de bij rechteroverschrijdingskans α_r behorende waarde van de $N(0,1)$ -verdeling is (tabel 1) en p_0 de getoetste waarde voorstelt. Deze formule geldt, zolang de uitkomst van het rechterlid $\leq n$ is; is deze $> n$, dan is de benaderde rechter kritieke zone leeg.

Lemma 15.1 doet vermoeden, dat x_r (eigenlijk wordt bedoeld: het rechterlid van (15.19)), bij gegeven n en α_r , een monotoon stijgende functie, van p_0 zal zijn. Daar het hier een benaderingsformule betreft is het wenselijk dit te controleren.

Daartoe beschouwen wij de afgeleide naar p (voor het gemak laten wij de indices 0 en α_r in deze berekening weg).

Deze is:

$$x'_r = n + \frac{1}{2}un(1-2p)/\sqrt{np(1-p)}.$$

Voor kleine p is deze inderdaad positief, maar voor p dicht bij 1 wordt hij negatief. Wij tonen nu aan dat dit pas het geval is voor zo grote p , dat daarvoor $x_r > n$ is. Daartoe berekenen wij het punt, waar $x'_r = 0$. Dit is het geval voor

$$\frac{1}{2}u(2p-1) = \sqrt{np(1-p)}.$$

Het rechterlid van deze vergelijking is, voor $0 < p < 1$, positief, het linkerlid dus ook. Dit betekent, dat de oplossing een waarde van p is, die $> \frac{1}{2}$ is.

Lossen wij de vergelijking nu op, dan verkrijgen wij als oplossing

$$p = \frac{1}{2}\left\{1 + \sqrt{\frac{n}{n+u^2}}\right\}; \quad 1-p = \frac{1}{2}\left\{1 - \sqrt{\frac{n}{n+u^2}}\right\}.$$

Invullen van deze waarden in (15.19) geeft na herleiding

$$x_r = \frac{1}{2}n\left\{1 + \sqrt{\frac{n+u^2}{n}}\right\} > n,$$

waarmede het gestelde bewezen is.

Wij kunnen nu dus (15.17) toepassen om de bij een waarde x behorende p_ℓ te berekenen. Dit geeft

$$(15.20) \quad x \approx np_\ell + u_{\alpha_\ell} \sqrt{np_\ell(1-p_\ell)} + \frac{1}{2},$$

waarin nu α_ℓ ($= \alpha_r$ uit (15.19)) de onbetrouwbaarheidsdrempel van p_ℓ voorstelt. Schrijven wij deze vergelijking als

$$x - \frac{1}{2} - np_\ell \approx u_{\alpha_\ell} \sqrt{np_\ell(1-p_\ell)},$$

dan zien wij, dat het rechterlid ≥ 0 is, zodat dus moet gelden

$$p_\ell \leq \frac{1}{n}(x - \frac{1}{2}).$$

Deze ongelijkheid geeft aan welke der beide oplossingen van de vierkants-

vergelijking, die door kwadrateren ontstaat, de juiste is. Deze blijkt te zijn

$$(15.21) \quad p_l \approx \frac{(x - \frac{1}{2}) + \frac{1}{2}u_{\alpha_l}^2 - \sqrt{u_{\alpha_l}^2 \left\{ \frac{(x - \frac{1}{2})(n - x + \frac{1}{2})}{n} + \frac{1}{4}u_{\alpha_l}^2 \right\}}}{n + u_{\alpha_l}^2}.$$

Op analoge wijze verkrijgt men voor de bovengrens

$$(15.22) \quad p_r \approx \frac{(x + \frac{1}{2}) + u_{\alpha_r}^2 + \sqrt{u_{\alpha_r}^2 \left\{ \frac{(x + \frac{1}{2})(n - x - \frac{1}{2})}{n} + \frac{1}{4}u_{\alpha_r}^2 \right\}}}{n + u_{\alpha_r}^2},$$

terwijl beide grenzen tezamen een betrouwbaarheidsinterval vormen met $\alpha = \alpha_r + \alpha_l$.

OPMERKING 15.5. Voor grote n en x wordt, door verwaarlozing van kleine termen in teller en noemer, formule (12.77) teruggekregen.

Zet men de zo verkregen grenzen voor $n = 100$ uit tegen x/n als abscis, dan wordt fig. 15.2 verkregen. Doet men dit voor verschillende waarden van n , dan verkrijgt men een bundel van dergelijke lijnen*).

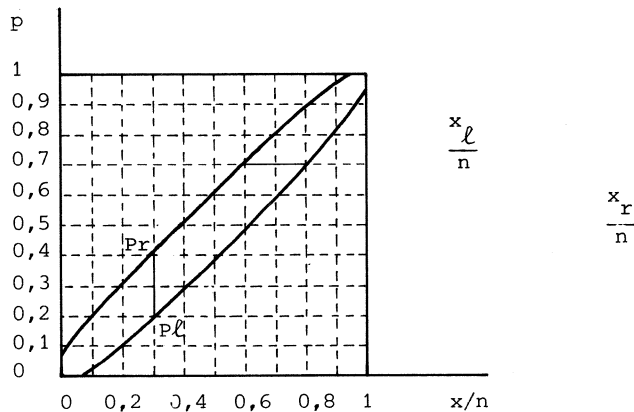


Fig. 15.2. Benaderde betrouwbaarheidsintervallen voor p

($n = 100$, $\alpha_l = \alpha_r = 0,025$)

*)

Zie bijv. W.J. DIXON & F.J. MASSEY Jr., *Introduction to statistical analysis*, 3rd ed., McGraw-Hill, N.Y.-Toronto-London 1969, table A-9 of L.N.H. BUNT, l.c. fig. 23.

Geven dus de verticale lijnen in deze figuur betrouwbaarheidsintervallen voor p aan (nl. p_ℓ en p_r als functie van x), men kan zich ook afvragen, wat de interpretatie is van de snijpunten der twee krommen met een horizontale lijn. De vergelijkingen der beide krommen worden gegeven door (15.20), waarvan het rechterlid (dus ook het linkerlid) x_r voorstelt, en een daaraan analoge uitdrukking voor x_ℓ . De snijpunten, bij constante p , zijn dus niet anders dan de kritieke waarden x_ℓ en x_r , behorende bij deze waarden van p als getoetse waarde.

Behalve uit de formule valt dit ook direct in te zien. Immers ieder punt buiten de "vis" komt overeen met een x en een p , die niet met elkaar in overeenstemming zijn: de waarde van p wordt, bij toetsing op grond van x , verworpen. De buiten de vis gelegen punten van een horizontale lijn (dus bij één p behorend) vormen dus tezamen alle waarden van x , die tot verwerping van deze waarde p leiden, dus juist de kritieke zones.

In de figuur komt duidelijk tot uiting, dat x_ℓ met p_r , en x_r met p_ℓ correspondeert.

Het oorspronkelijke door J. NEYMAN bij de afleiding van deze methode gegeven bewijs verliep juist omgekeerd^{*)}. Zijn redenering was: teken de kromme x_r en x_ℓ als functie van p_0 , de getoetste waarde. Er is één horizontale lijn (onbekend welke), die bij de werkelijke p behoort. Hierop komt dus het punt $(p, \underline{x})$ terecht. De kans, dat het links van (of in) x_ℓ valt is $\leq \alpha_\ell$ (als toetsings-onbetrouwbaarheid) en de kans, dat het rechts van x_r komt is $\leq \alpha_r$. Realiseert zich geen van deze beide mogelijkheden, dan komt het punt $(p, \underline{x})$ dus in de vis terecht (al weten wij niet waar) en deze kans is dus $\geq 1 - \alpha$. Wij voorspellen nu:

Het punt $(p, \underline{x})$ komt in de vis terecht, een voorspelling met kans $\geq 1 - \alpha$ op juistheid. Vervolgens wordt $\underline{x}$ waargenomen; dit geeft een waarde x . Het punt (p, x) ligt nu, volgens de voorspelling, in de vis, dus de bewering

Het punt (p, x) ligt in de vis heeft een betrouwbaarheid $\geq 1 - \alpha$. Maar deze bewering is equivalent met: p ligt tussen de snijpunten, die de verticale lijn door x met de beide krommen heeft. Dit zijn dus de betrouwbaarheidsgrenzen.

^{*)} In zijn allereerste publicatie hierover gaf NEYMAN een veel ingewikkelder bewijs, dat in een later artikel tot het hieronder gegevene werd herleid.

OPGAVE 15.1. In verband met de bij boekingen voor bungalows te volgen politiek wil een reisbureau inzicht hebben in de kans op afzeggen van een vroeg gemaakte reservering. Uit de gegevens van het vorige jaar blijkt, dat van 221 vroege boekingen 33 tenslotte niet zijn doorgegaan. Overwogen wordt in de winter en lente een zekere overboeking toe te laten. Daartoe wenst men te beschikken over een boven- en een ondergrens voor de kans op afzeggen van een reeds geboekte bungalow.

Bereken een betrouwbaarheidsinterval voor deze kans met tweezijdige onbetrouwbaarheid 0,01.

(Oplossing: $\alpha_l = \alpha_r = 0,005$, dus $u_{\alpha_l} = u_{\alpha_r} = 2,575$ (tabel 1). Invullen in (15.21) en (15.22) geeft $p_l = 0,096$ en $p_r = 0,224$.)

Eigenschappen van betrouwbaarheidsintervallen

DEFINITIE 15.2. Een betrouwbaarheidsgebied $\underline{G}$ heet *zuiver* met betrekking tot de parameter θ , indien voor iedere van de werkelijke waarde θ verschillende waarde θ' geldt:

$$(15.23) \quad P(\theta \in \underline{G}) \geq P(\theta' \in \underline{G}).$$

STELLING 15.3. *Indien een toets T , voor de hypothese $\theta = \theta_0$, berust op een continu verdeelde toetsingsgrootheid en voor iedere toegelaten waarde van θ_0 zuiver is, dan is het bijbehorende betrouwbaarheidsgebied voor θ ook zuiver.*

BEWIJS. De continuïteit van de toetsingsgrootheid heeft tot gevolg, dat onbetrouwbaarheid en onbetrouwbaarheidsdrempel gelijk zijn. Geef beide aan met α .

Wordt nu een onjuiste waarde θ' getoetst, dan is, wegens de zuiverheid van de toets, de kans op verwerping van θ' minstens gelijk aan α . De kans, dat θ' niet in $\underline{G}$ ligt, is dus $\geq \alpha$ (en in het algemeen uiteraard $> \alpha$).

Wordt echter de werkelijke waarde θ getoetst, dan is de kans op verwerping daarvan $\leq \alpha$. $\square$

OPMERKINGEN.

15.6. Voor tweezijdige toetsen is het begrip zuiverheid niet erg zinvol, omdat (vgl. blz. 148) het onderscheidingsvermogen in dat geval verkregen wordt door optelling van de kansen op twee geheel ongelijksoortige gebeur-

tenissen, nl. het trekken van de juiste conclusie enerzijds en het trekken van een falikant verkeerde anderzijds. Voor tweezijdige betrouwbaarheidsintervallen is het begrip zuiverheid echter wel zinvol.

15.7. Eénzijdige betrouwbaarheidsintervallen zijn nooit zuiver in bovengenoemde zin. Immers is, voor een ondergrens t_{ℓ} , de relatie $t_{\ell} < \theta$ juist, dan geldt ook $t_{\ell} < \theta'$ voor iedere $\theta' > \theta$. Dit betekent, dat in plaats van (15.23) geldt:

$$(15.24) \quad P(t_{\ell} < \theta) \leq P(t_{\ell} < \theta') \quad \text{voor iedere } \theta' > \theta,$$

hetgeen juist overeenkomt met onzuiverheid.

Naar de andere kant, nl. voor $\theta'' < \theta$, is een ondergrens echter altijd zuiver. Immers is $t_{\ell} < \theta''$, dan ook $t_{\ell} < \theta$ (daar $\theta'' < \theta$), dus geldt

$$(15.25) \quad P(t_{\ell} < \theta) \geq P(t_{\ell} < \theta'') \quad \text{voor iedere } \theta'' < \theta.$$

Analoge relaties gelden voor een bovengrens t_r . De al of niet zuiverheid van de bijbehorende toets is hierbij niet ter sprake gebracht, hetgeen misschien de indruk zou kunnen wekken, dat deze voor dit resultaat niet van belang is. Dit is echter slechts schijn. Immers het is gemakkelijk uit (15.24 en 25) de (éénzijdige) zuiverheid van de bijbehorende toets af te leiden. Dit betekent dus, dat de bij een éénzijdig betrouwbaarheidsinterval behorende toets altijd zuiver is; een onzuivere toets kan nooit tot een éénzijdig begrensde betrouwbaarheidsgebied leiden.

DEFINITIE 15.3. Een betrouwbaarheidsinterval voor θ heet *asymptotisch verdwijnend* als voor de (stochastische) lengte $L = t_r - t_{\ell}$ geldt:

$$(15.26) \quad \lim_{n \rightarrow \infty} P(L > \varepsilon) = 0 \quad \text{voor iedere } \varepsilon > 0.$$

STELLING 15.4. Wordt een betrouwbaarheidsondergrens t_{ℓ} afgeleid uit een éénzijdige toets T voor de hypothese $\theta \leq \theta_0$, die voor iedere θ_0 asymptotisch onderscheidend is met betrekking tot iedere $\theta' > \theta_0$, dan geldt:

$$(15.27) \quad \lim_{n \rightarrow \infty} P(t_{\ell} < \theta - \frac{1}{2}\varepsilon) = 0 \quad \text{voor iedere } \varepsilon > 0.$$

Hetzelfde geldt, *mutatis mutandis* voor t_r .

BEWIJS. $\underline{t}_\ell < \theta - \frac{1}{2}\epsilon$ betekent, dat de waarde $\theta - \frac{1}{2}\epsilon$ niet verworpen wordt. De kans hierop nadert echter voor $n \rightarrow \infty$ tot 0. $\square$

STELLING 15.5. *Wordt een tweezijdig betrouwbaarheidsinterval afgeleid uit een tweezijdige toets, samengesteld uit twee eenzijdige asymptotisch onderscheidende toetsen, dan is het betrouwbaarheidsinterval asymptotisch verdwijnend.*

BEWIJS. Volgens de vorige stelling geldt (15.27) en

$$(15.28) \quad \lim_{n \rightarrow \infty} P(\underline{t}_{-r} > \theta + \frac{1}{2}\epsilon) = 0 \quad \text{voor iedere } \epsilon > 0.$$

Daar voor ieder tweetal gebeurtenissen geldt:

$$P(A \vee B) \leq P(A) + P(B)$$

volgt hieruit, dat ook geldt:

$$\lim_{n \rightarrow \infty} P(\underline{t}_{-\ell} < \theta - \frac{1}{2}\epsilon \vee \underline{t}_{-r} > \theta + \frac{1}{2}\epsilon) = 0,$$

dus

$$\lim_{n \rightarrow \infty} P(\theta - \frac{1}{2}\epsilon \leq \underline{t}_{-\ell} < \underline{t}_{-r} \leq \theta + \frac{1}{2}\epsilon) = 1,$$

waaruit (15.26) volgt. $\square$

Heeft van twee toetsen T_1 en T_2 voor eenzelfde hypothese T_1 overal een groter onderscheidingsvermogen dan T_2 , dan verdient het bij T_1 behorende betrouwbaarheidsgebied $\underline{G}_1$ de voorkeur boven het bij T_2 behorende $\underline{G}_2$. Immers voor iedere $\theta' \neq \theta$ geldt dan

$$(15.30) \quad P(\theta' \in \underline{G}_1) < P(\theta' \in \underline{G}_2),$$

hetgeen betekent, dat $\underline{G}_1$ in de regel korter zal zijn dan $\underline{G}_2$.

VOORBEELD 15.2. Betrouwbaarheidsgrenzen voor de mediaan.

Laat van een grootheid $\underline{x}$ n onafhankelijke waarnemingen $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ beschikbaar zijn, terwijl gevraagd wordt een (éénzijdige of tweezijdige) betrouwbaarheidsinterval voor de mediaan x_M te bepalen. Wij veronderstellen - om complicaties te vermijden - dat $\underline{x}$ continu verdeeld is, zodat geen gelijke waarden onder de waarnemingen voorkomen (althans de kans daarop is 0). Voor x_M geldt dan:

$$(15.31) \quad P(\underline{x} < \underline{x}_M) = P(\underline{x} > \underline{x}_M) = \frac{1}{2}.$$

Wij rangschikken nu de waarnemingen naar opklimmende grootte en geven ze dan aan met $\underline{x}_{(1)}, \underline{x}_{(2)}, \dots, \underline{x}_{(n)}$, waarvoor dus geldt

$$(15.32) \quad P(\underline{x}_{(1)} < \underline{x}_{(2)} < \dots < \underline{x}_{(n)}) = 1.$$

Noem nu $\underline{a}$ het aantal der waarnemingen, die kleiner zijn dan een te toetsen waarde $\underline{x}_{M,0}$ voor de mediaan, dan bezit $\underline{a}$, onder de hypothese $\underline{x}_M = \underline{x}_{M,0}$, een binomiale verdeling met $p = \frac{1}{2}$, zodat voor toetsing van tabel 2, de tekentoets, gebruik gemaakt kan worden.

Laat a_ℓ de in die tabel gevonden (of eventueel de met behulp van de desbetreffende formule berekende) linker kritieke waarde zijn, behorende bij een onbetrouwbaarheidsdrempel α_ℓ . Dan wordt dus $\underline{x}_{M,0}$ verworpen indien $\underline{a}$ een waarde $\leq a_\ell$ aanneemt. Nu is echter $\underline{a} \leq a_\ell$ equivalent met: onder $\underline{x}_{M,0}$ liggen a_ℓ of minder waarnemingen. Of ook: $\underline{x}_{M,0} \leq \underline{x}_{(a_\ell+1)}$.

Alle waarden $\leq \underline{x}_{(a_\ell+1)}$ worden dus verworpen, zodat volgens stelling (15.1) geldt:

$$(15.33) \quad P(\underline{x}_{(a_\ell+1)} < \underline{x}_M) \geq 1 - \alpha_\ell.$$

De betrouwbaarheidsondergrens is dus de $(a_\ell+1)^e$ waarneming bij rangschikking naar grootte. Om redenen van symmetrie moet de betrouwbaarheidsboven-grens met $\alpha_r = \alpha_\ell$, op dezelfde wijze gevonden worden, maar nu van "rechts" naar "links" werkende. In plaats van de $(a_\ell+1)^e$ van onderen af moet nu dus de $(a_\ell+1)^e$ van boven af genomen worden. Deze draagt het nummer $n-a_\ell$. (N.B.: niet $n-a_\ell-1!$), zodat wij verkrijgen:

$$(15.34) \quad P(\underline{x}_M < \underline{x}_{(n-a_\ell)}) \geq 1 - \alpha_\ell.$$

Een tweezijdige betrouwbaarheidsinterval wordt verkregen door (15.33) en (15.34) te combineren. De notatie kan vereenvoudigd worden als wij $\alpha_\ell = \alpha_r = \frac{1}{2}\alpha$ stellen. Geven wij de linker-kritieke waarde van de teken-toets voor dat geval aan met r , zodat de rechter-kritieke waarde $n-r$ wordt, dan geldt:

$$(15.35) \quad P(\underline{x}_{(r+1)} < \underline{x}_M < \underline{x}_{(n-r)}) \geq 1 - \alpha,$$

terwijl de beide grenzen ieder afzonderlijk $\frac{1}{2}\alpha$ als onbetrouwbaarheids-

drempel hebben.

OPMERKING 15.8. Hoewel wij verondersteld hebben, dat $\underline{x}$ continu verdeeld is, is toch de onbetrouwbaarheid hier niet gelijk aan de onbetrouwbaarheidsdrempel, omdat de toetsingsgrootte discreet verdeeld is.

Is $\underline{x}$ discreet verdeeld, dan treden complicaties op, die wij hier buiten beschouwing laten.

OPGAVE 15.2. Het tingealte van 20 aselekt uit een grondlaag genomen monsters bedraagt, uitgedrukt in de een of andere eenheid:

21,4; 13,7; 14,0; 18,2; 7,7; 16,3; 27,4; 17,4; 23,1; 5,6; 22,0; 16,9; 27,9; 22,8; 10,8; 23,5; 13,5; 19,0; 17,8; 18,0.

Bepaal een tweezijdig betrouwbaarheidsinterval voor de mediaan van het tingealte, met onbetrouwbaarheidsdrempel 0,05.

(Oplossing: Volgens tabel 2 is $r = 5$. Bij rangschikking naar grootte blijken de grenzen te zijn: 14,0 en 22,0.)

VOORBEELD 15.3. Betrouwbaarheidsgrenzen voor n

Indien van een reeks van n onafhankelijke experimenten, met voor ieder kans p op succes, de waarde van p bekend is, maar n niet, kan men door directe toepassing van stelling 15.1 betrouwbaarheidsgrenzen voor deze onbekende n bepalen. De binomiale toets wordt daarbij dan niet gebruikt als toets voor hypothesen van de vorm $p = p_0$, maar, bij de bekende p , voor hypothesen van de vorm $n = n_0$.

OPGAVE 15.3. Bij een reeks van n onafhankelijke experimenten, ieder met kans $\frac{1}{2}$ op succes, worden 20 successen verkregen. Het aantal (n) is onbekend.

- Gevraagd wordt een tweezijdig betrouwbaarheidsinterval voor n te bepalen met $\alpha = 0,05$.
 - Los de opgave ook op als $p = \frac{2}{3}$ is.
 - Leid (met behulp van de normale benadering zonder continuïteitscorrectie) een algemene formule voor dit betrouwbaarheidsinterval af.
- (Opgave 21 uit het Sommenboekje; oplossing op blz. 63-65.)

VOORBEELD 15.4. Betrouwbaarheidsgrenzen voor een symmetriepunt

Indien de grootte $\underline{x}$, waarvan n onafhankelijke waarnemingen beschikbaar zijn, symmetrisch verdeeld is om een onbekend punt ξ (dat tevens $= x_M$ en

= μ is), kan de symmetrietoets van Wilcoxon gebruikt worden om een betrouwbaarheidsinterval voor ξ te bepalen.

De procedure is eenvoudig, maar eist, indien men geen vaardigheid in de toepassing bezit, nogal wat rekenwerk. Het komt eenvoudig hierop neer, dat men in plaats van de waarnemingen zelf, de getallen

$$(15.36) \quad x_1 - \Delta, x_2 - \Delta, \dots, x_n - \Delta$$

beschouwt en daarop, voor verschillende waarden van Δ , de tweezijdige symmetrietoets toepast. Er zijn twee grenswaarden voor Δ , Δ_r en Δ_l , waarbij juist verworpen moet worden, terwijl dit voor intermediaire waarden niet het geval is. Dit zijn dan de betrouwbaarheidsgrenzen voor ξ .

De procedure kan ook éénzijdig worden uitgevoerd. De kritieke waarden van de toets kunnen aan tabel 4 of aan de desbetreffende formules worden ontleend.

OPGAVE 15.4. De gegevens van opgave 15.2 zijn ontleend aan een symmetrische verdeling. Bepaal een tweezijdige betrouwbaarheidsinterval voor het symmetriepunt, met $\alpha = 0,05$.

OPLOSSING.

Δ	T	
15,0	109	} $n = 20$, rechter kritieke waarde = 106
15,1	108	
15,2	105	
21	-122	} $n = 20$, linker kritieke waarde = -106
20,7	-110	
20,6	-104	

Het betrouwbaarheidsinterval is dus

$$15,1 < \xi < 20,7.$$

VOORBEELD 15.5. Betrouwbaarheidsinterval voor de verschuiving van één verdeling ten opzichte van een andere

Indien twee stochastische grootheden $\underline{x}$ en $\underline{y}$ op een verschuiving na dezelfde verdeling bezitten, d.w.z. als

(15.37) $\underline{x} + \Delta$ en $\underline{y}$

voor een onbekende waarde van Δ dezelfde verdeling bezitten, dan kan, op grond van twee onafhankelijke steekproeven $\underline{x}_1, \dots, \underline{x}_m$ en $\underline{y}_1, \dots, \underline{y}_n$ een betrouwbaarheidsinterval voor de "verschuiving" Δ bepaald worden.

Dit geschiedt geheel analoog aan het vorige voorbeeld, maar nu met behulp van de toets van Wilcoxon voor twee steekproeven. De hypothese $\Delta = \Delta_0$ kan, met behulp van die toets, voor iedere Δ_0 getoetst worden en de niet verworpen waarden vormen tezamen het betrouwbaarheidsgebied.

De uitvoering is in dit geval eenvoudiger dan in het geval van de symmetrietoets. Door de beide steekproeven ieder op een strookje millimeterpapier uit te zetten en deze langs elkaar te schuiven, zijn de kritieke verschuivingen (waarbij nog juist niet of net wel verworpen wordt) gemakkelijk te bepalen. Bij verschuiving van het ene strookje ten opzichte van het andere verandert de waarde van de toetsingsgrootte met 2 wanneer een x-streepje een y-streepje passeert.

OPGAVE 15.5.

Bepaal uit de gegevens van opgave 14.2.3 een tweezijdig betrouwbaarheidsinterval, met $\alpha = 0,05$, voor de verschuiving van de verdelingen ten opzichte van elkaar.

OPLOSSING. $m = 10$, $n = 13$, $\alpha = 0,05$; de kritieke waarden van W zijn dus 66 en 194.

Zonder verschuiving ($\Delta = 0$):

x _____
y _____ $W = 114$

Met verschuiving van x over 0,93 naar rechts:

x _____
y _____ $W = 193$

Met verschuiving van x over 0,69 naar links:

x _____
y _____ $W = 67$

Het interval is:

$$-0,93 \leq \Delta \leq 0,69.$$

Afleiding van betrouwbaarheidsgrenzen met behulp van de schattingstheorie

Een andere methode om betrouwbaarheidsintervallen af te leiden berust op de schattingstheorie. Het principe van deze methode kan als volgt uiteengezet worden.

Laat $\underline{t}$ een schatter van de onbekende parameter θ zijn met $E\underline{t} = \theta$ en $\sigma^2(\underline{t}) = \sigma^2$ onafhankelijk van θ . Laat verder de verdeling van $\underline{t}$, op de onbekendheid van θ na, bekend zijn, bijv. normaal. Dan geldt:

$$(15.38) \quad P(\theta - u_{\frac{1}{2}\alpha} \sigma < \underline{t} < \theta + u_{\frac{1}{2}\alpha} \sigma) = 1 - \alpha.$$

Het interval $(\theta - u_{\frac{1}{2}\alpha} \sigma, \theta + u_{\frac{1}{2}\alpha} \sigma)$ heet een *voorspellingsinterval* voor $\underline{t}$; de voorspelling, dat $\underline{t}$ in dit interval zal liggen heeft een kans $1 - \alpha$ juist te zijn.

Dit interval laat zich echter gemakkelijk omrekenen tot een betrouwbaarheidsinterval voor θ . Immers (15.38) is equivalent met

$$(15.39) \quad P(\underline{t} - u_{\frac{1}{2}\alpha} \sigma < \theta < \underline{t} + u_{\frac{1}{2}\alpha} \sigma) = 1 - \alpha$$

en het interval $(\underline{t} - u_{\frac{1}{2}\alpha} \sigma, \underline{t} + u_{\frac{1}{2}\alpha} \sigma)$ is nu een betrouwbaarheidsinterval met onbetrouwbaarheid α .

Op dit principe, zij het soms in gewijzigde vorm, berust de afleiding van vele betrouwbaarheidsintervallen.

Volgens opmerking 15.5 behoort bij zo'n interval ook weer een toets, zodat er geen essentieel verschil met de in het voorafgaande behandelde methode is.

Een moeilijkheid is echter, dat de verdeling van $\underline{t}$ vaak niet zo eenvoudig van θ afhangt, dat een vergelijking van de vorm (15.38) opgeschreven kan worden.

Beschouwen wij bijv. het geval van een binomiale grootte $\underline{x}$ met parameters n en p , waarvoor dus geldt: $E\underline{x} = np$ en $\sigma^2(\underline{x}) = npq$, dan hangen verwachting en spreiding beide van n en van p af. Zowel indien bij bekende n een betrouwbaarheidsinterval voor p bepaald moet worden (voorbeeld 15.1) als wanneer p bekend is en een betrouwbaarheidsinterval voor n gezocht wordt (voorbeeld 15.3), leidt dit tot complicaties, die gewoonlijk

omzeild worden op de in hoofdstuk XII beschreven wijze: vervanging van σ door een schatter daarvan, zonder rekening te houden met het stochastische karakter daarvan (resultaat: formule (12.77)). Deze moeilijkheid treedt echter niet op indien men van de toetsingstheorie uitgaat.

In nog sterkere mate geldt dit voor de voorbeelden 15.2, 4 en 5. In die gevallen berust de toets op een toetsingsgrootheid, die geen schatter van de onbekende parameter is, zodat directe toepassing van de schattingstheorie niet mogelijk is. En hoewel men in alle drie genoemde gevallen wel een schatter voor de onbekende parameter kan aangeven, is de verdeling van deze schatter of wel onbekend, ofwel afhankelijk van de vorm van de verdeling der waarnemingen ofwel beide, zodat het met behulp van de toetsingstheorie opgeloste probleem uitgaande van de schattingstheorie niet opgelost is of zelfs niet kan worden.

Anderzijds zijn er echter gevallen, waarin de schattingstheorie even snel als de toetsingstheorie tot een oplossing (en dan ook tot dezelfde oplossing) voert.

VOORBEELD 15.6. Betrouwbaarheidsgrenzen voor de verwachting van een normale verdeling met onbekende spreiding

Indien $\underline{x} \sim N(\mu, \sigma)$ -verdeeld is, en $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_n$ vormen een steekproef uit deze verdeling, dan bezit volgens stelling 14.2.1 de grootheid

$$(15.40) \quad \underline{t} = \frac{\bar{\underline{x}} - \mu}{\underline{s}} \sqrt{n}$$

een Student-verdeling met $n - 1$ vrijheidsgraden.

In tabel 2 kan men nu de waarde $t_{\frac{1}{2}\alpha}$ opzoeken, die een tweezijdige overschrijdskans α bezit en daarvoor geldt dus

$$(15.41) \quad P(|\underline{t}| < t_{\frac{1}{2}\alpha}) = 1 - \alpha,$$

of, na enige herleiding

$$(15.42) \quad P(\bar{\underline{x}} - t_{\frac{1}{2}\alpha} \frac{\underline{s}}{\sqrt{n}} < \mu < \bar{\underline{x}} + t_{\frac{1}{2}\alpha} \frac{\underline{s}}{\sqrt{n}}) = 1 - \alpha.$$

Dit is een tweezijdig betrouwbaarheidsinterval met onbetrouwbaarheid α , ieder der betrouwbaarheidsgrenzen, die de eindpunten van dit interval zijn, bezit afzonderlijk onbetrouwbaarheid $\frac{1}{2}\alpha$.

OPGAVE 15.6. Bereken, in de veronderstelling van normaliteit, een tweezijdig betrouwbaarheidsinterval voor μ , met onbetrouwbaarheid 0,05, uit de gegevens van opgave 15.2.

OPLOSSING. $\bar{x} = \frac{1}{n} \sum x_i = 17,85$; $s^2 = \frac{1}{n-1} \sum (x_i - \bar{x})^2 = 34,99$; $s = 5,8$;
 $t_{\frac{1}{2}\alpha} = t_{0,025} = 2,093$ ($v = 19$). De grenzen worden: 15,1 en 20,6.

VOORBEELD 15.7. Betrouwbaarheidsgrenzen voor het verschil van de verwachting van twee normale verdelingen met gelijke spreiding

Indien $\underline{x}$ en $\underline{y}$ beide normaal verdeeld zijn, met verwachtingen μ_1 resp. μ_2 en met dezelfde spreiding σ , terwijl $\underline{x}_1, \underline{x}_2, \dots, \underline{x}_m$ en $\underline{y}_1, \underline{y}_2, \dots, \underline{y}_n$ twee onafhankelijke steekproeven uit deze verdelingen zijn, dan heeft volgens stelling 14.2.2 de grootheid

$$(15.43) \quad \underline{t} = \frac{\bar{x} - \bar{y} - (\mu_1 - \mu_2)}{\sqrt{\frac{s_1^2 + s_2^2}{m+n}}} \sqrt{\frac{mn(m+n-2)}{m+n}}$$

een Student-verdeling met $m + n - 2$ vrijheidsgraden.

Derhalve geldt dan weer (15.41), waarbij $t_{\frac{1}{2}\alpha}$ in de juiste regel van tabel 2 moet worden opgezocht.

Nu valt deze vergelijking te herleiden tot

$$(15.44) \quad P\left(\bar{x} - \bar{y} - t_{\frac{1}{2}\alpha} \sqrt{\frac{(s_1^2 + s_2^2)}{mn(m+n-2)}} \frac{m+n}{mn(m+n-2)} < \mu_1 - \mu_2 < \bar{x} - \bar{y} + t_{\frac{1}{2}\alpha} \sqrt{\frac{(s_1^2 + s_2^2)}{mn(m+n-2)}} \frac{m+n}{mn(m+n-2)}\right) = 1 - \alpha.$$

OPGAVE 15.7. Bepaal, in de veronderstelling van normaliteit, uit de gegevens van 14.2.3 een tweezijdig betrouwbaarheidsinterval, met $\alpha = 0,05$, voor $\mu_1 - \mu_2$.

(Opgave 80b van het Sommenboekje; oplossing op blz. 119:

$$-0,755 < \mu_1 - \mu_2 < 0,609.)$$

VOORBEELD 15.8. Betrouwbaarheidsgrenzen voor het verschil van twee onbekende kansen

Een toets voor de gelijkheid van 2 kansen ($H_0: p_1 = p_2$) is in het begin van hoofdstuk XIV behandeld. Deze toets leent zich echter niet voor het

toetsen van hypothesen van de vorm $p_1 - \Delta = p_2$, zodat hieraan geen betrouwbaarheidsinterval voor $\Delta = p_1 - p_2$ te ontlenen valt. Een dergelijk interval moet dus op andere wijze afgeleid worden en hiervoor is geen andere methode bekend dan de aan het eind van hoofdstuk XII reeds beschreven methode "voor grote steekproeven".

Beschikken wij over 2 onafhankelijke reeksen van n_1 resp. n_2 proeven, waarvan de ene betrekking heeft op p_1 en de andere op p_2 , en zijn de aantallen successen x_1 resp. x_2 , dan zijn deze grootheden binomiaal verdeeld en onafhankelijk van elkaar.

Als schatter voor p_1 resp. p_2 kunnen x_1/n_1 resp. x_2/n_2 dienen (vgl. (12.31)) en als schatter voor $p_1 - p_2$ gebruiken wij daarom de grootheid

$$(15.45) \quad \underline{d} = \frac{x_1}{n_1} - \frac{x_2}{n_2}.$$

Voor de verwachting en de variantie van $\underline{d}$ geldt:

$$(15.46) \quad E\underline{d} = p_1 - p_2; \quad \sigma^2(\underline{d}) = \frac{p_1 q_1}{n_1} + \frac{p_2 q_2}{n_2}.$$

(De lezer bewijze dit.)

Zijn n_1 en n_2 niet te klein, dan zijn x_1 en x_2 en derhalve ook $\underline{d}$ bij benadering normaal verdeeld, zodat geldt:

$$(15.47) \quad P(|\underline{d} - (p_1 - p_2)| < u_{\frac{1}{2}\alpha} \sigma(\underline{d})) \approx 1 - \alpha.$$

Hieruit volgt een benaderd betrouwbaarheidsinterval, als men $\sigma(\underline{d})$ door een schatting van die grootheid vervangt. Daartoe vult men gewoonlijk in de formule voor $\sigma^2(\underline{d})$ (15.46) de schatters x_1/n_1 resp. x_2/n_2 in voor p_1 resp. p_2 , hetgeen leidt tot de schatter:

$$(15.48) \quad \underline{s}^2 = \frac{x_1(n_1 - x_1)}{n_1^3} + \frac{x_2(n_2 - x_2)}{n_2^3}.$$

Soms gebruikt men ook noemers $n_1^2(n_1 - 1)$ en $n_2^2(n_2 - 1)$, om de schatter zuiver te krijgen (vgl. (12.76)), maar de betekenis van deze correctie is beperkt, omdat dan wel voor σ^2 , maar niet voor σ , een zuivere schatter verkregen wordt, terwijl bovendien voor grote n de invloed van de correctie gering is.

Het (tweezijdige) betrouwbaarheidsinterval, al of niet met deze correctie, wordt dus:

$$(15.49) \quad P\left(\frac{x_1}{n_1} - \frac{x_2}{n_2} - u_{\frac{1}{2}\alpha} s < p_1 - p_2 < \frac{x_1}{n_1} - \frac{x_2}{n_2} + u_{\frac{1}{2}\alpha} s\right) \approx 1 - \alpha.$$

Beide grenzen afzonderlijk bezitten als (benaderde) onbetrouwbaarheidsdrempel $\frac{1}{2}\alpha$.

OPGAVE 15.8. Bereken een tweezijdig betrouwbaarheidsinterval met $\alpha = 0,01$, voor het verschil $p_1 - p_2$ van de ziektekans van selderieplanten op grond van de gegevens van voorbeeld 14.1.1.

OPLOSSING.

$$\frac{x_1}{n_1} = 0,2; \quad \frac{x_2}{n_2} = 0,46; \quad s^2 = \frac{20 \cdot 80}{10^6} + \frac{46 \cdot 54}{10^6} = 0,004084;$$

$$s = 0,0639.$$

Het betrouwbaarheidsinterval wordt (met $u_{\frac{1}{2}\alpha} = 2,575$):

$$-0,42 < p_1 - p_2 < -0,10.$$

OPMERKING 15.9. De bij dit betrouwbaarheidsinterval behorende toets stelt ons in staat de hypothese $p_1 - p_2 = d_0$, met gegeven waarde van d_0 , bij benadering te toetsen door na te gaan of d_0 in het betrouwbaarheidsinterval ligt of niet. De methode der 2×2 -tabel (par. 14.1) geldt alleen voor $d_0 = 0$, maar is voor dat geval exact uitvoerbaar voor kleine n_1 en n_2 , terwijl voor grote n_1 en n_2 de bijbehorende benadering (zie blz. 159) beter is dan de hier beschrevene. Voor $d_0 = 0$ is dus de laatstgenoemde methode te prefereren.

Hiermede zijn een aantal der ganbare schattings- en toetsingsmethoden beschreven. Een veel groter aantal is buiten beschouwing gelaten. Het in deze cursus bestreken gebied bevat bovendien slechts een zeer klein gedeelte van de momenteel beschikbare statistische technieken. Lezers, die zich verder in de mathematische statistiek wensen te verdiepen kunnen de in de inleiding vermelde leer- en handboeken raadplegen.

APPENDIX

§1. Normale verdeling.

Bewijs van de formule

$$(A.1.1) \quad \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} e^{-\frac{1}{2}\frac{(x-\mu)^2}{\sigma^2}} dx = 1.$$

Geven wij het linkerlid aan door L en substitueren wij $u = \frac{x-\mu}{\sigma}$, dan komt er, daar $dx = \sigma du$ is,

$$L = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} e^{-\frac{1}{2}u^2} du.$$

Derhalve is

$$L^2 = \frac{1}{2\pi} \iint_{-\infty}^{+\infty} e^{-\frac{1}{2}(u^2+v^2)} dudv$$

en overgang op poolcoördinaten r en ϕ geeft, daar $r^2 = u^2 + v^2$ en $dudv = r dr d\phi$

$$L^2 = \frac{1}{2\pi} \int_0^{2\pi} d\phi \int_0^{\infty} r e^{-\frac{1}{2}r^2} dr = \int_0^{\infty} e^{-\frac{1}{2}r^2} d(\frac{1}{2}r^2) = 1.$$

Afleiding van de momenten (formule (9.9))

Volgens (9.4) is

$$\mu_1 = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} x e^{-\frac{1}{2}\frac{(x-\mu)^2}{\sigma^2}} dx.$$

Substitueer $u = \frac{x-\mu}{\sigma}$; $x = u\sigma + \mu$, $dx = \sigma du$:

$$\begin{aligned}\mu_1 &= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} (u\sigma + \mu) e^{-\frac{1}{2}u^2} du = \\ &= \frac{\sigma}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} u e^{-\frac{1}{2}u^2} du + \frac{\mu}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} e^{-\frac{1}{2}u^2} du.\end{aligned}$$

De eerste term valt weg wegens symmetrie van het argument onder de integraal en de tweede term is volgens (A.1.1) gelijk aan μ . Dus is

$$(A.1.2) \quad \mu_1 = E\bar{x} = \mu.$$

Verder is

$$\begin{aligned}\mu_2 &= \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} x^2 e^{-\frac{1}{2}\frac{(x-\mu)^2}{\sigma^2}} dx = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} (\sigma u + \mu)^2 e^{-\frac{1}{2}u^2} du = \\ &= \frac{\sigma^2}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} u^2 e^{-\frac{1}{2}u^2} du + \mu^2,\end{aligned}$$

daar nu de middelste term wegvalt. Nu is, met partieel integreren,

$$\int_{-\infty}^{+\infty} u^2 e^{-\frac{1}{2}u^2} du = \left[-u e^{-\frac{1}{2}u^2} \right]_{-\infty}^{+\infty} + \int_{-\infty}^{+\infty} e^{-\frac{1}{2}u^2} du = \sqrt{2\pi},$$

hetgeen na invullen geeft

$$(A.1.3) \quad \mu_2 = E\bar{x}^2 = \sigma^2 + \mu^2.$$

§2. Exponentiële verdeling

Als $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_n$ onderling onafhankelijk verdeeld zijn volgens dezelfde exponentiële verdeling

$$(A.2.1) \quad P(\bar{x}_i \leq x_i) = 1 - e^{-\lambda x_i} \quad x_i \geq 0,$$

dan is

$$(A.2.2) \quad E \frac{1}{\bar{x}} = \frac{n}{n-1} \lambda \quad (\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i).$$

BEWIJS. Als

$$(A.2.3) \quad \underline{y} \stackrel{\text{def}}{=} \sum_{i=1}^n x_i,$$

dan is de verdelingsdichtheid $f(y)$ van $\underline{y}$

$$(A.2.4) \quad f(y) = \frac{\lambda^n}{(n-1)!} y^{n-1} e^{-\lambda y}.$$

Dit kan als volgt met behulp van volledige inductie bewezen worden. Uit

(A.2.1) volgt voor de verdelingsdichtheid $g(x_i)$ van x_i

$$(A.2.5) \quad g(x_i) = \lambda e^{-\lambda x_i},$$

dus voor $n = 2$ is (zie (11.5))

$$(A.2.6) \quad f(y) = \int_0^y g(x) g(y-x) dx = \lambda^2 \int_0^y e^{-\lambda x} e^{-\lambda(y-x)} dx = \\ = \lambda^2 \int_0^y e^{-\lambda y} dx = \lambda^2 y e^{-\lambda y}.$$

Stel nu dat (A.2.4) juist is voor $n - 1$, dan wordt de verdelingsdichtheid voor de som van n grootheden x_i

$$(A.2.7) \quad \int_0^y \frac{\lambda^{n-1}}{(n-2)!} x^{n-2} e^{-\lambda x} \cdot \lambda e^{-\lambda(y-x)} dx = \frac{\lambda^n}{(n-2)!} e^{-\lambda y} \int_0^y x^{n-2} dx = \\ = \frac{\lambda^n}{(n-1)!} y^{n-1} e^{-\lambda y}.$$

Uit (A.2.4) volgt nu

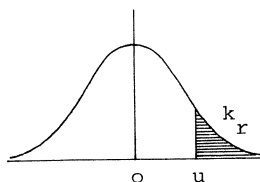
$$(A.2.8) \quad E \frac{1}{\underline{y}} = \frac{\lambda^n}{(n-1)!} \int_0^\infty y^{n-2} e^{-\lambda y} dy = \frac{\lambda^n}{(n-1)!} \frac{(n-2)!}{\lambda^{n-1}} = \frac{\lambda}{n-1},$$

daar (zie (A.2.7))

$$\frac{\lambda^n}{(n-1)!} \int_0^{\infty} y^{n-1} e^{-\lambda y} dy = 1.$$

Dus is

$$(A.2.9) \quad E \frac{1}{\underline{x}} = nE \frac{1}{\underline{y}} = \frac{n}{n-1} \lambda. \quad \square$$



TABEL 1

N(0,1)-verdeling

Waarden van $k_r \cdot 10^4$ voor $u = 0,00(0,01)3,49$.

$$k_r = \frac{1}{\sqrt{2\pi}} \int_u^{\infty} e^{-\frac{1}{2}v^2} dv.$$

u	0	1	2	3	4	5	6	7	8	9
0,0	5000	4960	4920	4880	4840	4801	4761	4721	4681	4641
0,1	4602	4562	4522	4483	4443	4404	4364	4325	4286	4247
0,2	4207	4168	4129	4090	4052	4013	3974	3936	3897	3859
0,3	3821	3783	3745	3707	3669	3632	3594	3557	3520	3483
0,4	3446	3409	3372	3336	3300	3264	3228	3192	3156	3121
0,5	3085	3050	3015	2981	2946	2912	2877	2843	2810	2776
0,6	2743	2709	2676	2643	2611	2578	2546	2514	2483	2451
0,7	2420	2389	2358	2327	2296	2266	2236	2206	2177	2148
0,8	2119	2090	2061	2033	2005	1977	1949	1922	1894	1867
0,9	1841	1814	1788	1762	1736	1711	1685	1660	1635	1611
1,0	1587	1562	1539	1515	1492	1469	1446	1423	1401	1379
1,1	1357	1335	1314	1292	1271	1251	1230	1210	1190	1170
1,2	1151	1131	1112	1093	1075	1056	1038	1020	1003	0985
1,3	0968	0951	0934	0918	0901	0885	0869	0853	0838	0823
1,4	0808	0793	0778	0764	0749	0735	0721	0708	0694	0681
1,5	0668	0655	0643	0630	0618	0606	0594	0582	0571	0559
1,6	0548	0537	0526	0516	0505	0495	0485	0475	0465	0455
1,7	0446	0436	0427	0418	0409	0401	0392	0384	0375	0367
1,8	0359	0351	0344	0336	0329	0322	0314	0307	0301	0294
1,9	0287	0281	0274	0268	0262	0256	0250	0244	0239	0233
2,0	0228	0222	0217	0212	0207	0202	0197	0192	0188	0183
2,1	0179	0174	0170	0166	0162	0158	0154	0150	0146	0143
2,2	0139	0136	0132	0129	0125	0122	0119	0116	0113	0110
2,3	0107	0104	0102	0099	0096	0094	0091	0089	0087	0084
2,4	0082	0080	0078	0075	0073	0071	0069	0068	0066	0064
2,5	0062	0060	0059	0057	0055	0054	0052	0051	0049	0048
2,6	0047	0045	0044	0043	0041	0040	0039	0038	0037	0036
2,7	0035	0034	0033	0032	0031	0030	0029	0028	0027	0026
2,8	0026	0025	0024	0023	0023	0022	0021	0021	0020	0019
2,9	0019	0018	0018	0017	0016	0016	0015	0015	0014	0014
3,0	0013	0013	0013	0012	0012	0011	0011	0011	0010	0010
3,1	0010	0009	0009	0009	0008	0008	0008	0008	0007	0007
3,2	0007	0007	0006	0006	0006	0006	0006	0005	0005	0005
3,3	0005	0005	0005	0004	0004	0004	0004	0004	0004	0003
3,4	0003	0003	0003	0003	0003	0003	0003	0003	0003	0002

Een 5 betekent, dat bij afronding naar het dichtsbijgelegen oneven cijfer afgerond moet worden.

Een 5 wordt naar het dichtsbijzijnde even cijfer afgerond.

Tabel 2

waarden van t_v bij tweezijdige overschrijdingskansen k van de Student-verdeling

$k \backslash v$	0,9	0,8	0,7	0,6	0,5	0,4	0,3	0,2	0,1	0,05	0,02	0,01
1	0,158	0,325	0,510	0,727	1,000	1,376	1,963	3,078	6,314	12,706	31,821	63,657
2	0,142	0,289	0,445	0,617	0,816	1,061	1,386	1,886	2,920	4,303	6,965	9,925
3	0,137	0,277	0,424	0,584	0,765	0,978	1,250	1,638	2,353	3,182	4,541	5,841
4	0,134	0,271	0,414	0,569	0,741	0,941	1,190	1,533	2,132	2,776	3,747	4,604
5	0,132	0,267	0,408	0,559	0,727	0,920	1,156	1,476	2,015	2,571	3,365	4,032
6	0,131	0,265	0,404	0,553	0,718	0,906	1,134	1,440	1,943	2,447	3,143	3,707
7	0,130	0,263	0,402	0,549	0,711	0,896	1,119	1,415	1,895	2,365	2,998	3,499
8	0,130	0,262	0,399	0,546	0,706	0,889	1,108	1,397	1,860	2,306	2,896	3,355
9	0,129	0,261	0,398	0,543	0,703	0,883	1,100	1,383	1,833	2,262	2,821	3,250
10	0,129	0,260	0,397	0,542	0,700	0,879	1,093	1,372	1,812	2,228	2,764	3,169
11	0,129	0,260	0,396	0,540	0,697	0,876	1,088	1,363	1,796	2,201	2,718	3,106
12	0,128	0,259	0,395	0,539	0,695	0,873	1,083	1,356	1,782	2,179	2,681	3,055
13	0,128	0,259	0,394	0,538	0,694	0,870	1,079	1,350	1,771	2,160	2,650	3,012
14	0,128	0,258	0,393	0,537	0,692	0,868	1,076	1,345	1,761	2,145	2,624	2,977
15	0,128	0,258	0,393	0,536	0,691	0,866	1,074	1,341	1,753	2,131	2,602	2,947
16	0,128	0,258	0,392	0,535	0,690	0,865	1,071	1,337	1,746	2,120	2,583	2,921
17	0,128	0,257	0,392	0,534	0,689	0,863	1,069	1,333	1,740	2,110	2,567	2,898
18	0,127	0,257	0,392	0,534	0,688	0,862	1,067	1,330	1,734	2,101	2,552	2,878
19	0,127	0,257	0,391	0,533	0,688	0,861	1,066	1,328	1,729	2,093	2,539	2,861
20	0,127	0,257	0,391	0,533	0,687	0,860	1,064	1,325	1,725	2,086	2,528	2,845
21	0,127	0,257	0,391	0,532	0,686	0,859	1,063	1,323	1,721	2,080	2,518	2,831
22	0,127	0,256	0,390	0,532	0,686	0,858	1,061	1,321	1,717	2,074	2,508	2,819
23	0,127	0,256	0,390	0,532	0,685	0,858	1,060	1,319	1,714	2,069	2,500	2,807
24	0,127	0,256	0,390	0,531	0,685	0,857	1,059	1,318	1,711	2,064	2,492	2,797
25	0,127	0,256	0,390	0,531	0,684	0,856	1,058	1,316	1,708	2,060	2,485	2,787
26	0,127	0,256	0,390	0,531	0,684	0,856	1,058	1,315	1,706	2,056	2,479	2,779
27	0,127	0,256	0,389	0,531	0,684	0,855	1,057	1,314	1,703	2,052	2,473	2,771
28	0,127	0,256	0,389	0,530	0,683	0,855	1,056	1,313	1,701	2,048	2,467	2,763
29	0,127	0,256	0,389	0,530	0,683	0,854	1,055	1,311	1,699	2,045	2,462	2,756
30	0,127	0,256	0,389	0,530	0,683	0,854	1,055	1,310	1,697	2,042	2,457	2,750
∞	0,12566	0,25335	0,38532	0,52440	0,67449	0,84162	1,03643	1,28155	1,64485	1,95996	2,32634	2,57582

Linker kritieke waarden van de tweezijdige tekentoets^{*)}

n	Onbetrouwbaarheidsdrempel (α)				n	Onbetrouwbaarheidsdrempel (α)			
	0,01	0,02	0,05	0,10		0,01	0,02	0,05	0,10
1	-	-	-	-	51	15	16	18	19
2	-	-	-	-	52	16	17	18	19
3	-	-	-	-	53	16	17	18	20
4	-	-	-	-	54	17	18	19	20
5	-	-	-	0	55	17	18	19	20
6	-	-	0	0	56	17	18	20	21
7	-	0	0	0	57	18	19	20	21
8	0	0	0	1	58	18	19	21	22
9	0	0	1	1	59	19	20	21	22
10	0	0	1	1	60	19	20	21	23
11	0	1	1	2	61	20	20	22	23
12	1	1	2	2	62	20	21	22	24
13	1	1	2	3	63	20	21	23	24
14	1	2	2	3	64	21	22	23	24
15	2	2	3	3	65	21	22	24	25
16	2	2	3	4	66	22	23	24	25
17	2	3	4	4	67	22	23	25	26
18	3	3	4	5	68	22	23	25	26
19	3	4	4	5	69	23	24	25	27
20	3	4	5	5	70	23	24	26	27
21	4	4	5	6	71	24	25	26	28
22	4	5	5	6	72	24	25	27	28
23	4	5	6	7	73	25	26	27	28
24	5	5	6	7	74	25	26	28	29
25	5	6	7	7	75	25	26	28	29
26	6	6	7	8	76	26	27	28	30
27	6	7	7	8	77	26	27	29	30
28	6	7	8	9	78	27	28	29	31
29	7	7	8	9	79	27	28	30	31
30	7	8	9	10	80	28	29	30	32
31	7	8	9	10	81	28	29	31	32
32	8	8	9	10	82	28	30	31	33
33	8	9	10	11	83	29	30	32	33
34	9	9	10	11	84	29	30	32	33
35	9	10	11	12	85	30	31	32	34
36	9	10	11	12	86	30	31	33	34
37	10	10	12	13	87	31	32	33	35
38	10	11	12	13	88	31	32	34	35
39	11	11	12	13	89	31	33	34	36
40	11	12	13	14	90	32	33	35	36
41	11	12	13	14	91	32	33	35	37
42	12	13	14	15	92	33	34	36	37
43	12	13	14	15	93	33	34	36	38
44	13	13	15	16	94	34	35	37	38
45	13	14	15	16	95	34	35	37	38
46	13	14	15	16	96	34	36	37	39
47	14	15	16	17	97	35	36	38	39
48	14	15	16	17	98	35	37	38	40
49	15	15	17	18	99	36	37	39	40
50	15	16	17	18	100	36	37	39	41

*) " - " betekent dat voor de betreffende waarden van n en α voor geen enkele waarde van de toetsingsgrootte de tweezijdige overschrijdingskans $\leq \alpha$ is.

Tabel 4

Rechter kritieke waarden van de tweezijdige symmetrie-
toets van Wilcoxon *)

n	Onbetrouwbaarheidsdrempel (α)			
	0,01	0,02	0,05	0,10
4	-	-	-	-
5	-	-	-	15
6	-	-	21	17
7	-	28	24	22
8	36	34	30	26
9	43	39	35	29
10	49	45	39	35
11	56	52	46	40
12	64	60	52	44
13	73	67	57	49
14	81	75	63	55
15	90	82	70	60
16	98	90	78	66
17	107	99	85	71
18	117	107	91	77
19	126	116	98	84
20	136	124	106	90
21	151	135	115	97
22	161	145	123	103
23	172	154	130	110
24	182	164	140	118
25	193	175	147	125
26	205	185	157	131
27	216	196	166	138
28	228	206	174	146
29	241	217	183	155
30	253	229	193	161

*) " - " betekent dat voor de betreffende waarden van n en α voor geen enkele waarde van T de tweezijdige overschrijdingskans $\leq \alpha$ is. Voor $n \leq 20$ zijn de kritieke waarden exact berekend, daarboven met behulp van de normale benadering.

Tabel 5a *)

Kritieke waarden van de tweezijdige toets voor twee steekproeven van Wilcoxon voor $m \leq n \leq 15$ en $\alpha = 0,02$

$m \backslash n$	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
2	-	-	-	-	-	-	-	-	-	-	-	-	52	56	60
3	-	-	-	-	-	-	42	48	52	58	64	68	74	80	84
4	-	-	-	-	40	46	54	60	66	74	80	86	94	100	106
5	-	-	-	0	48	56	64	72	80	88	96	104	112	120	128
6	-	-	-	2	4	66	76	84	94	104	114	122	132	142	150
7	-	-	0	2	6	8	12	86	98	108	118	130	140	150	162
8	-	-	0	4	8	12	14	18	110	122	134	146	158	168	180
9	-	-	2	6	10	14	18	22	28	134	148	162	174	188	200
10	-	-	2	6	12	16	22	26	32	38	162	176	192	206	220
11	-	-	2	8	14	18	24	30	36	44	50	192	208	224	240
12	-	-	4	10	16	22	28	34	42	48	56	62	226	242	260
13	-	0	4	12	18	24	32	40	46	54	62	70	260	278	296
14	-	0	4	12	20	26	34	44	52	60	68	76	86	298	318
15	-	0	6	14	22	30	38	48	56	66	74	84	94	102	338

*)

"-" betekent dat voor de betreffende waarden van m en n voor geen enkele waarde van w de tweezijdige overschrijdingskans $\leq \alpha$ is.

Tabel 5b
Kritieke waarden van de tweezijdige toets voor twee steekproeven van Wilcoxon voor $m \leq n \leq 15$ en $\alpha = 0,05$ *)

$\begin{smallmatrix} m \\ n \end{smallmatrix}$	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
2	-	-	-	-	-	-	-	32	36	40	44	46	50	54	58
3	-	-	-	-	30	34	40	44	50	54	60	64	70	74	80
4	-	-	-	32	38	44	50	56	64	70	76	82	88	94	100
5	-	-	0	2	4	54	60	68	76	84	92	98	106	114	122
6	-	-	2	4	6	62	72	80	88	98	106	116	124	134	142
7	-	-	2	6	10	12	16	20	26	34	40	46	52	60	68
8	-	0	4	8	12	16	20	24	30	34	40	46	52	60	68
9	-	0	4	8	14	20	24	28	32	36	40	44	48	52	56
10	-	0	6	10	16	22	28	32	36	40	44	48	52	56	60
11	-	0	6	12	18	26	32	36	40	44	48	52	56	60	64
12	-	2	8	14	22	28	36	40	44	48	52	56	60	64	68
13	-	2	8	16	24	32	40	44	48	52	56	60	64	68	72
14	-	2	10	18	26	34	44	48	52	56	60	64	68	72	76
15	-	2	10	20	28	38	48	52	56	60	64	68	72	76	80

*) Zie voetnoot bij tabel 5a.

Tabel 5c

Kritieke waarden van de tweezijdige toets voor twee steekproeven van Wilcoxon voor $m \leq n \leq 15$ en $\alpha = 0,10$ *)

$\begin{smallmatrix} m \\ n \end{smallmatrix}$	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
1	-	-	-	-	-	-	-	-	-	-	-	-	-	-	-
2	-	-	-	-	20	24	28	30	34	38	42	44	48	50	54
3	-	-	18	24	28	32	38	42	46	52	56	62	66	70	76
4	-	-	0	30	26	42	48	54	60	66	72	78	84	90	96
5	-	0	2	4	8	42	50	64	72	78	86	94	100	108	114
6	-	0	4	6	10	14	58	68	76	84	92	100	118	126	134
7	-	0	4	8	12	16	22	76	86	96	106	116	134	144	154
8	-	2	6	10	16	20	26	30	98	108	120	140	152	162	174
9	-	2	8	12	18	24	30	36	42	120	132	144	168	180	192
10	-	2	8	14	22	28	34	40	48	54	146	158	186	198	212
11	-	2	10	16	24	32	38	46	54	62	68	174	202	216	230
12	-	4	10	18	26	34	42	52	60	68	76	84	204	234	250
13	-	4	12	20	30	38	48	56	66	74	84	94	236	252	268
14	-	6	14	22	32	42	52	62	72	82	92	102	112	122	288
15	-	6	14	24	36	46	56	66	78	88	100	110	122	132	306
															144

*) Zie voetnoot bij tabel 5a.

REGISTER

- Aannemelijkheidsfunctie*, 115
- ACTON, F.S., 6
- Afhankelijkheid*: van eventualiteiten, 33, 34, 36;
van stochastische grootheden, 67. Zie ook *Onafhankelijkheid*
- ANDERSON, T.W., 7
- Aselect*, 37, 38, 39, 40, 41, 44, 52, 55
- Asymptotisch*: doeltreffend, 133; normaal verdeeld, 96, 97, 103;
onderscheidend, 149, 165, 174, 199, 201; raak, 106, 107, 110,
112, 132; verdwijnend, 212
- BAHADUR, R.R., 200
- BAILLE, A., 8
- BARRA, J.P., 5, 8
- BAYES: stelling van, 20
- BENARD, A., 181, 184
- Betrouwbaarheids*: bovengrens, 201; coëfficiënt, 129, 203;
gebied, 204, 211; ondergrens, 201
- Betrouwbaarheidsgrenzen*: afleiding uit schattingstheorie, 218;
afleiding uit toetsingstheorie, 203
- Betrouwbaarheidsinterval*, 129, 201, 204; asymptotisch ver-
dwindend, 212; voor de mediaan, 213; voor de verschuiving
van één verdeling t.o.v. een andere, 216; voor de verwachting
van een normale verdeling met onbekende spreiding, 219; voor
een onbekende kans, 204, 205, 206; voor een onbekende kans,
benaderd, 134, 207, 208, 209, 210; voor een symmetriepunt, 215;
voor het verschil van de verwachtingen van twee normale verde-
lingen met gelijke spreiding, 220; voor het verschil van twee
onbekende kansen, benaderd, 220; voor n , 215; zuiver, 211
- BEYER, W.H., 8
- BIENAYMÉ-TCHEBYCHEFF: ongelijkheid van, 85
- BISHOP, Y.M.M., 7
- BOES, D., 5
- BRADLEY, J.V., 5, 201
- BRUIJN, N.G. DE, 44
- BUNT, L.N.H., 206, 209

- CAMERON, J.M., 172
- CHUNG, K.L., 8
- COCHRAN, W.G., 5, 7
- Continuïteitscorrectie*, 151, 185
- Correlatiecoëfficiënt*, 83, 93
- Covariantie*, 82, 83, 93
- COX, D.R., 6
- CRAMÉR, H., 96, 97
- CRAIG, A.T., 5
- DANIEL, C., 6
- DANTZIG, D. VAN, 29, 37, 192
- Dichotomie*, 70; dubbele, 156, 195
- DIXON, W.J., 5, 209
- Doeltreffendheid*, 200
- DRAPER, M.R., 6
- EEDEN, C. VAN, 181, 184, 187, 195
- Efficiency*: asymptotisch relatieve, 200; Bahadur, 200; Hodges & Lehmann, 200; Pitman, 200; relatieve, 199, 200
- ETTINGER, J. VAN, 111
- Eventualiteit*, 29
- FELLER, W., 7
- FERGUSON, T.S., 5
- FERMAT, 10
- FIENBERG, S.E., 7
- FINETTI, B. DE, 7
- FISHER, R.A., 8, 114
- Fout*: van de eerste soort, 143; van de tweede soort, 146
- Frequentiequotiënt*, 11, 14, 26, 35; voorwaardelijke, 14, 31
- Frequentieverdeling*, 118
- Gemiddelde*, 86, 91, 105; schatting voor, 107; voorlopig, 91, 169
- Gestandaardiseerde grootheid*, 79, 85
- GIBBONS, J.D., 6
- GOSSET, W.S., zie STUDENT

GRAYBILL, F.A., 5, 6

GREENWOOD, J.A., 9

HABERMAN, S.J., 7

HÁJEK, J., 6, 192, 193, 199, 200, 201

HARTLEY, H.O., 8, 9

HEMELRIJK, J., 4, 8, 52, 199

HODGES, J.L., 4, 200

HOGG, R.V., 5

HOLLAND, P.W., 7

HOLLANDER, M., 5, 193

HUYGENS, 24

Hypothesen: alternatieve, 140; enkelvoudige, 138; samengestelde, 139;
te toetsen (getoetste, nul), 139; toegelaten, 138; verworpen, 140

INSTITUTE OF MATHEMATICAL STATISTICS, 8

JOHN, P.W.M., 6

JOHNSON, N.L., 8

JONGE, H. DE, 4

Kans: frequentie interpretatie, 25; voorwaardelijke, 30, 36, 40

Kansrekening, 2, 10, 11, 14, 24, 36, 45, 46

Kansverdeling, 47, 48; continue, 49, 50; discrete, 48, 49; gemengde, 51;
marginale, 55, 56; n-dimensionale, 52, 53, 54; n-dimensionale
continue, 54; n-dimensionale discrete, 52; simultane, 52;
voorwaardelijke, 55, 57, 58, 59

KEMPTHORNE, O., 6

KENDALL, M.G., 5, 7, 8, 193

Kleinste kwadraten: theorie der, 122, 127

KOLMOGOROV, A., 35

KOTZ, S., 8

KRES, H., 8

Kritieke: waarde, 142; zone, 141

KRUSKAL, W.H., 4, 192

LAPLACE: kansdefinitie van, 24, 34, 37, 39, 41, 66;

stelling van, 97, 98

LEHMANN, E.L., 4, 5, 6, 200

Limietstelling: centrale, 90, 95, 96, 103

LINDGREN, B.W., 5

LOÈVE, M., 7

MANN, H.B., 192, 193

MASSEY, F.J., 5, 209

Mediaan, 130

MENDEL, 137

MENDENHALL, W., 6

MÉRÉ, CHEVALLIER, DE, 10, 25

MILLER, R.G., 7

MISES, R. VON, 25

Model, 1, 2, 13, 25, 28, 29, 34, 37, 46, 47, 55, 65, 66

MOIVRE, DE: stelling van, 97, 98

Moment: eerste, 70; eerste gereduceerde, 72; k-de, 70; k-de gereduceerde, 72; relatie tussen gereduceerde en gewone, 75; tweede, 70; tweede gereduceerde, 72

MOOD, A.M., 5

MORRISON, D.F., 7

MOSTELLER, F., 4

NEVEU, J., 7

NEYMAN, J., 210

NOETHER, G.E., 4, 5

NÖLLE, G., 5

Onafhankelijk, 32, 36, 44

Onafhankelijkheid: bij inrichting van experimenten, 65, 66;

en marginale verwachting, 73; en voorwaardelijke kansen, 62;

en voorwaardelijke verdelingen, 65; praktische interpretatie,

33, 34; relatie met correlatie van normaal verdeelde grootheden, 66;

relatie met covariantie, 82; relatie met de verwachting van een product van stochastische grootheden, 74; van eventualiteiten, 32, 33, 34; van stochastische grootheden, 61; van continu verdeelde stochastische grootheden, 63, 64; van discreet verdeelde stochastische grootheden, 62, 63;

Onbetrouwbaarheid: linker, 145; rechter, 145; van een interval of gebied, 130, 201, 206; van een toets, 143

Onbetrouwbaarheidsdrempel: van een toets, 143; van een betrouwbaarheidsinterval of -gebied, 204

Onderscheidingsvermogen, 145, 146; asymptotisch, 149, 199; linker, 146; lokaal, 199; rechter, 146; van de binomiale toets, 146, 147; van de Student toets voor één steekproef, 196, 197; van de Student toets voor twee steekproeven, 197; van de tekentoets, 180, 186, 198; van de toets voor de 2×2 -tabel, 161; van de Wilcoxon symmetrie-toets, 186, 199; van de Wilcoxon toets voor twee steekproeven, 199

Overschrijdingskans, 150, 152, 158, 159, 168, 169, 173, 174, 179, 185, 193, 194; linker, 150; rechter, 150; tweezijdig, 150

OWEN, D.B., 8

Parameter, 105, 106, 136, 138

Parameterruimte, 105, 138

Parameter vrij, 165

PARZEN, E., 7

PASCAL, 10

PEARSON, E.S., 8, 166

PITMAN, E.J.G., 200

Populatie, 112

PRESS, S.J., 7

PURI, M.L., 6

QUÉTELET, A., 10

RAHMAN, N.A., 8

RAY, D., 7

RAND CORPORATION, 8

RAO, C.R., 5

Regressie: coëfficiënt, 134; lineaire, 122, 134; lijn, 125; waarde, 125

- Residu*, 127
- ROBERTS, H.V., 4
- ROMIG, H.G., 144, 147
- RIJKOORT, P.J., 193
- SEARLE, S.R., 6
- SEN, P.K., 6
- Schatter*, 106, 107, 109, 112; asymptotisch doeltreffende, 133;
asymptotisch rake, 106, 110, 112; meest nauwkeurig zuivere,
128; nauwkeurigheid van een, 128; zuivere, 106;
- Schatter, meest aannemelijke*, 115; voor binomiale p , 117; voor
gemiddelde en variantie van de normale verdeling, 115; voor
negatief binomiale p , 118; voor parameter van de exponentiële
verdeling, 118, 119, 120, 121; voor regressie coëfficiënten,
124, 125, 126; voor standaard deviatie van de normale verde-
ling, 116
- Schatter voor*: binomiale p , 112; de mediaan, 132; de variantie van
de normale verdeling, 106; het gemiddelde van de normale
verdeling, 105; negatief binomiale p , 113
- Schatting*, 106
- Schattingstheorie*, 104, 118
- SCHEFFÉ, H., 6
- SCHMETTERER, L., 5
- SIDÁK, Z., 6
- Simulatie*, 199
- SMID, L.J., 193
- SMITH, H., 6
- SNEDECOR, G.W., 5
- SOM, K.K., 7
- Spreiding*, zie *standaardafwijking*
- Spreidingsbreedte*, 175
- Standaardafwijking*, 72; zie ook *variantie*
- Stochastisch*, 13; punt, 52
- Stochastische*: convergentie, 86, 87; onafhankelijkheid van eventualitei-
ten, 32; onafhankelijkheid van stochastische grootheden, 61
- Stochastische grootheden*, 47; functies van, 69, 90; onafhankelijk
verdeelde, 61; simultane kansverdeling van een aantal, 52
- STOKER, D.J., 192

STUART, A., 7, 8

STUDENT, 166, 174, 179, 185, 196, 219

SÜSSKIND, F.W.J., 10

SVERDRUP, E., 5

Systeem: exclusief, 17, 30; kategorisch, 17, 41

TATSUOKA, M.M., 6

TIMM, N.H., 7

Toets: asymptotisch onderscheidende, 149; asymptotisch onderscheidingsvermogen van een, 199; keuze tussen een één- of tweezijdige, 155; linkséénzijdige, 142; lokaal onderscheidingsvermogen van een, 199; onbetrouwbaarheid van een, 143; onderscheidingsvermogen van een, 145; parameter vrije, 165; rechtséénzijdige, 142; symmetrisch tweezijdige, 145; tweezijdige, 142; verdelingsvrije, 165; zuivere, 149

Toets voor een binomiale kans, 136, 139, 140, 141, 142, 143, 144, 145, 146, 147, 148; eigenschappen van, 149, 150; normale benadering van, 151, 152, 153, 154

Toets voor één steekproef van Student, 168; asymptotisch onderscheidend, 174, 201; onderscheidingsvermogen, 180, 196, 197, 198; toepassing op verschillen van paren waarnemingen, 171; vervuld zijn der voorwaarden, 168, 170

Toets voor één steekproef, de Tekentoets, 177; knopen, 179; lokaal- en asymptotisch onderscheidingsvermogen, 201; onderscheidingsvermogen, 197, 198; Pitman-efficiency, 200; toepassing op verschillen van paren waarnemingen, 179; vervuld zijn der voorwaarden, 178; voor hypothesen m.b.t. de mediaan, 178

Toets voor één steekproef van Wilcoxon, symmetrie toets, 180; asymptotisch onderscheidend, 201; exacte verdeling toetsingsgrootheid, 184; knopen, 181, 182, 185; lokaal- en asymptotisch onderscheidingsvermogen, 200, 201; normale benadering, 181, 182, 184; onderscheidingsvermogen, 186, 199; Pitman-efficiency, 200; toepassing op verschillen van paren waarnemingen, 181

Toets voor gelijkheid van 2 binomiale kansen, 156, 157, 158; asymptotisch onderscheidend, 161; eenzijdige toetsen zuiver, 161; normale benadering, 159; onderscheidingsvermogen, 161; verschil der frakties en chi-kwadraat toets voor dezelfde hypothese, 166

Toets voor gelijkheid van 2 verdelingen, mediaan toets, 162, 163, 164, 165; asymptotisch onderscheidend, 165; verschil der frakties en chi-kwadraat toets voor dezelfde hypothese, 166

Toets voor onafhankelijkheid van 2 kenmerken, 161, 162; chi-kwadraat toets voor dezelfde hypothese, 166

Toets voor 2 steekproeven van Student, 169; asymptotisch onderscheidend, 201; onderscheidingsvermogen, 197; vervuld zijnder voorwaarden, 168, 175

Toets voor 2 steekproeven van Wilcoxon, 187; asymptotisch onderscheidend, 201; exacte verdeling toetsingsgrootheid, 192, 193; knopen, 188, 190, 192, 195; lokaal- en asymptotisch onderscheidingsvermogen, 200, 201; normale benadering, 188, 192; onderscheidingsvermogen, 199; Pitman-efficiency, 200; relatie met toets voor 2×2 -tabel, 195

Toetsen: van Student, 166; verdelingsvrije, 177; voor de 2×2 -tabel, 156

Toetsingsgrootheid, 140

Toetsingstheorie, 136

TUCKER, H., 8

Variantie, 72; van een lineaire combinatie, 83; van een schatter, 128; van een som van onafhankelijke grootheden, 76; van een steekproef, 107; van een stochastische grootheid, 72, 75, 76; van het gemiddelde van een steekproef, 128; van toetsingsgrootheid symmetrie toets van Wilcoxon, 181, 182, 183; van toetsingsgrootheid toets voor gelijkheid van 2 kansen, 157; van toetsingsgrootheid 2 steekproeven toets van Wilcoxon, 188, 189, 190, 191

Verdeling: asymptotisch normale, 158, 182, 188; binomiale, 41, 48, 53, 77, 112, 117, 133, 141, 204, 206; chi-kwadraat, 166; continue, 49, 54, 63, 69; n-dimensionale, 54, 61, 72; discrete, 48, 52, 55, 62, 67, 69; exponentiële, 118; gestandaardiseerde normale, 79, 80, 91, 100, 101; homogene, 51, 53, 71, 75; hypergeometrische, 41, 48, 78, 112, 157; marginale, 55, 64, 65, 73, 74; negatief binomiale, 43, 49, 67, 70, 77, 113, 118; normale, 50, 51, 57, 59, 65, 66, 73, 74, 78, 90, 92, 93; simultane, 52, 54, 56, 57, 63, 66, 67, 72, 80, 82, 90, 115; tweedimensionale normale, 54, 55, 57, 59, 65, 73, 74, 82, 83; tweedimensionale homogene, 55, 75; van Student, 166, 167, 168, 219; voorwaardelijke, 57, 58, 59, 65, 73, 157, 183

Verdelingsdichtheid: bij continue kansverdelingen, 49; voorwaardelijke, bij continue kansverdelingen, 58, 59

Verdelingsfunctie, 48; bij continue verdelingen, 49; bij continue marginale verdelingen, 56; bij continue voorwaardelijke verdelingen, 57; 58; 59; bij discrete verdelingen, 48; bij discrete marginale verdelingen, 55; bij discrete voorwaardelijke verdelingen, 57; zie ook *kansverdeling*

Verwachting, 69, 73; als lineaire operator, 75; marginale, 73; van een lineaire combinatie van stochastische grootheden, 75; van een product van stochastische grootheden, 74; voorwaardelijke, 73, 74

Voorspellingsinterval, 218

Waarschijnlijkheidsrekening, zie *kansrekening*

WABEKE, D., 8, 67, 187

WALLIS, W.A., 4

WEMELSFELDER, P.J., 118

Wet der grote aantallen: experimentele, 10, 11; theoretische, 87, 88, 89

WHITNEY, D.R., 192, 193

WILCOXON, F., 180, 181, 184, 187, 188, 193, 194, 195, 196, 199, 200, 201, 217

WINER, B.J., 6

WISHART, J., 166

WITTING H., 5

WOLFE, D.A., 6

WOOD, F.S., 6

YATES, F., 8

YAUDEN, W.J., 172

ZACKS, S., 5

Zuiverheid: van een betrouwbaarheidsgebied, 211; van een schatter, 106; van een toets, 149

UITGAVEN IN DE SERIE MC SYLLABUS

Onderstaande uitgaven zijn verkrijgbaar bij het Mathematisch Centrum,
2e Boerhaavestraat 49 te Amsterdam-1005, tel. 020-947272.

-
- | | |
|----------|---|
| MCS 1.1 | F. GÖBEL & J. VAN DE LUNE, <i>Leergang Besliskunde, deel 1: Wiskundige basiskennis</i> , 1965. ISBN 90 6196 014 2. |
| MCS 1.2 | J. HEMELRIJK & J. KRIENS, <i>Leergang Besliskunde, deel 2: Kansberekening</i> , 1965. ISBN 90 6196 015 0. |
| MCS 1.3 | J. HEMELRIJK & J. KRIENS, <i>Leergang Besliskunde, deel 3: Statistiek</i> , 1966. ISBN 90 6196 016 9. |
| MCS 1.4 | G. DE LEVE & W. MOLENAAR, <i>Leergang Besliskunde, deel 4: Markovketens en wachttijden</i> , 1966. ISBN 90 6196 017 7. |
| MCS 1.5 | J. KRIENS & G. DE LEVE, <i>Leergang Besliskunde, deel 5: Inleiding tot de mathematische besliskunde</i> , 1966. ISBN 90 6196 018 5. |
| MCS 1.6a | B. DORHOUT & J. KRIENS, <i>Leergang Besliskunde, deel 6a: Wiskundige programmering 1</i> , 1968. ISBN 90 6196 032 0. |
| MCS 1.6b | B. DORHOUT, J. KRIENS & J.TH. VAN LIESHOUT, <i>Leergang Besliskunde, deel 6b: Wiskundige programmering 2</i> , 1977. ISBN 90 6196 150 5. |
| MCS 1.7a | G. DE LEVE, <i>Leergang Besliskunde, deel 7a: Dynamische programmering 1</i> , 1968. ISBN 90 6196 033 9. |
| MCS 1.7b | G. DE LEVE & H.C. TIJMS, <i>Leergang Besliskunde, deel 7b: Dynamische programmering 2</i> , 1970. ISBN 90 6196 055 x. |
| MCS 1.7c | G. DE LEVE & H.C. TIJMS, <i>Leergang Besliskunde, deel 7c: Dynamische programmering 3</i> , 1971. ISBN 90 6196 066 5. |
| MCS 1.8 | J. KRIENS, F. GÖBEL & W. MOLENAAR, <i>Leergang Besliskunde, deel 8: Minimaxmethode, netwerkplanning, simulatie</i> , 1968. ISBN 90 6196 034 7. |
| MCS 2.1 | G.J.R. FÖRCH, P.J. VAN DER HOUWEN & R.P. VAN DE RIET, <i>Colloquium Stabiliteit van differentieschema's, deel 1</i> , 1967. ISBN 90 6196 023 1. |
| MCS 2.2 | L. DEKKER, T.J. DEKKER, P.J. VAN DER HOUWEN & M.N. SPIJKER, <i>Colloquium Stabiliteit van differentieschema's, deel 2</i> , 1968. ISBN 90 6196 035 5. |
| MCS 3.1 | H.A. LAUWERIER, <i>Randwaardeproblemen, deel 1</i> , 1967. ISBN 90 6196 024 x. |
| MCS 3.2 | H.A. LAUWERIER, <i>Randwaardeproblemen, deel 2</i> , 1968. ISBN 90 6196 036 3. |
| MCS 3.3 | H.A. LAUWERIER, <i>Randwaardeproblemen, deel 3</i> , 1968. ISBN 90 6196 043 6. |
| MCS 4 | H.A. LAUWERIER, <i>Representaties van groepen</i> , 1968. ISBN 90 6196 037 1. |

- MCS 5 J.H. VAN LINT, J.J. SEIDEL & P.C. BAAYEN, *Colloquium Discrete wiskunde*, 1968.
ISBN 90 6196 044 4.
- MCS 6 K.K. KOKSMA, *Cursus ALGOL 60*, 1969. ISBN 90 6196 045 2.
- MCS 7.1 *Colloquium Moderne rekenmachines, deel 1*, 1969. ISBN 90 6196 046 0.
- MCS 7.2 *Colloquium Moderne rekenmachines, deel 2*, 1969. ISBN 90 6196 047 9.
- MCS 8 H. BAVINCK & J. GRASMAN, *Relaxatietrillingen*, 1969.
ISBN 90 6196 056 8.
- MCS 9.1 T.M.T. COOLEN, G.J.R. FÖRCH, E.M. DE JAGER & H.G.J. PIJLS, *Elliptische differentiaalvergelijkingen, deel 1*, 1970.
ISBN 90 6196 048 7.
- MCS 9.2 W.P. VAN DEN BRINK, T.M.T. COOLEN, B. DIJKHUIS, P.P.N. DE GROEN, P.J. VAN DER HOUWEN, E.M. DE JAGER, N.M. TEMME & R.J. DE VOGELAERE, *Colloquium Elliptische differentiaalvergelijkingen, deel 2*, 1970.
ISBN 90 6196 049 5.
- MCS 10 J. FABIUS & W.R. VAN ZWET, *Grondbegrippen van de waarschijnlijkheidsrekening*, 1970. ISBN 90 6196 057 6.
- MCS 11 H. BART, M.A. KAASHOEK, H.G.J. PIJLS, W.J. DE SCHIPPER & J. DE VRIES, *Colloquium Halfalgebra's en positieve operatoren*, 1971.
ISBN 90 6196 067 3.
- MCS 12 T.J. DEKKER, *Numerieke algebra*, 1971. ISBN 90 6196 068 1.
- MCS 13 F.E.J. KRUSEMAN ARETZ, *Programmeren voor rekenautomaten; De MC ALGOL 60 vertaler voor de EL X8*, 1971. ISBN 90 6196 069 X.
- MCS 14 H. BAVINCK, W. GAUTSCHI & G.M. WILLEMS, *Colloquium Approximatie-theorie*, 1971. ISBN 90 6196 070 3.
- MCS 15.1 T.J. DEKKER, P.W. HEMKER & P.J. VAN DER HOUWEN, *Colloquium Stijve differentiaalvergelijkingen, deel 1*, 1972. ISBN 90 6196 078 9.
- MCS 15.2 P.A. BEENTJES, K. DEKKER, H.C. HEMKER, S.P.N. VAN KAMPEN & G.M. WILLEMS, *Colloquium Stijve differentiaalvergelijkingen, deel 2*, 1973. ISBN 90 6196 079 7.
- MCS 15.3 P.A. BEENTJES, K. DEKKER, P.W. HEMKER & M. VAN VELDHUIZEN, *Colloquium Stijve differentiaalvergelijkingen, deel 3*, 1975.
ISBN 90 6196 118 1.
- MCS 16.1 L. GEURTS, *Cursus Programmeren, deel 1: De elementen van het programmeren*, 1973. ISBN 90 6196 080 0.
- MCS 16.2 L. GEURTS, *Cursus Programmeren, deel 2: De programmeertaal ALGOL 60*, 1973. ISBN 90 6196 087 8.
- MCS 17.1 P.S. STOBBE, *Lineaire algebra, deel 1*, 1974. ISBN 90 6196 090 8.
- MCS 17.2 P.S. STOBBE, *Lineaire algebra, deel 2*, 1974. ISBN 90 6196 091 6.
- MCS 17.3 N.M. TEMME, *Lineaire algebra, deel 3*, 1976. ISBN 90 6196 123 8.
- MCS 18 F. VAN DER BLIJ, H. FREUDENTHAL, J.J. DE IONGH, J.J. SEIDEL & A. VAN WIJNGAARDEN, *Een kwart eeuw wiskunde 1946-1971, Syllabus van de Vakantiecursus 1971*, 1974. ISBN 90 6196 092 4.
- MCS 19 A. HORDIJK, R. POTHARST & J.Th. RUNNENBURG, *Optimaal stoppen van Markovketens*, 1974. ISBN 90 6196 093 2.

- MCS 20 T.M.T. COOLEN, P.W. HEMKER, P.J. VAN DER HOUWEN & E. SLAGT, *ALGOL 60 procedures voor begin- en randwaardeproblemen*, 1976. ISBN 90 6196 094 0.
- MCS 21 J.W. DE BAKKER (red.), *Colloquium Programmacorrectheid*, 1975. ISBN 90 6196 103 3.
- MCS 22 R. HELMERS, F.H. RUYMGAART, M.C.A. VAN ZUYLEN & J. OOSTERHOFF, *Asymptotische methoden in de toetsingstheorie; Toepassingen van naburigheid*, 1976. ISBN 90 6196 104 1.
- MCS 23.1 J.W. DE ROEVER (red.), *Colloquium Onderwerpen uit de biomathemica, deel 1*, 1976. ISBN 90 6196 105 X.
- * MCS 23.2 J.W. DE ROEVER (red.), *Colloquium Onderwerpen uit de biomathemica, deel 2*, 1976. ISBN 90 6196 115 7.
- MCS 24.1 P.J. VAN DER HOUWEN, *Numerieke integratie van differentiaalvergelijkingen, deel 1: Eenstapsmethoden*, 1974. ISBN 90 6196 106 8.
- MCS 25 *Colloquium Structuur van programmeertalen*, 1976. ISBN 90 6196 116 5.
- MCS 26.1 N.M. TEMME (ed.), *Nonlinear analysis, volume 1*, 1976. ISBN 90 6196 117 3.
- MCS 26.2 N.M. TEMME (ed.), *Nonlinear analysis, volume 2*, 1976. ISBN 90 6196 121 1.
- MCS 27 M. BAKKER, P.W. HEMKER, P.J. VAN DER HOUWEN, S.J. POLAK & M. VAN VELDHUIZEN, *Colloquium Discretiseringsmethoden*, 1976. ISBN 90 6196 124 6.
- MCS 28 O. DIEKMANN, N.M. TEMME (EDS), *Nonlinear Diffusion Problems*, 1976. ISBN 90 6196 126 2.
- MCS 29.1 J.C.P. BUS (red.), *Colloquium Numerieke programmatuur, deel 1A, deel 1B*, 1976. ISBN 90 6196 128 9.
- * MCS 29.2 H.J.J. TE RIELE (red.), *Colloquium Numerieke programmatuur, deel 2*, 1976. ISBN 144 0.
- * MCS 30 P. GROENEBOOM, R. HELMERS, J. OOSTERHOFF & R. POTTHARST, *Efficiency begrippen in de statistiek*, 1977. ISBN 90 6196 149 1.
- MCS 31 J.H. VAN LINT (red.), *Inleiding in de coderingstheorie*, 1976. ISBN 90 6196 136 X.
- MCS 32 L. GEURTS (red.), *Colloquium Bedrijfssystemen*, 1976. ISBN 90 6196 137 8.
- MCS 33 P.J. VAN DER HOUWEN, *Differentieschema's voor de berekening van waterstanden in zeeën en rivieren*, ISBN 90 6196 138 6.
- MCS 34 J. HEMELRIJK, *Oriënterende cursus mathematische statistiek*, ISBN 90 6196 139 4.
- * MCS 35 P.J.W. TEN HAGEN (red.), *Colloquium Computer Graphics*, 1977. ISBN 90 6196 142 4.
- * MCS 36 J.M. AARTS, J. DE VRIES (red.), *Colloquium Topologische Dynamische Systemen*, 1977. ISBN 90 6196 143 2.

De met een * gemerkte uitgaven moeten nog verschijnen.