

STAtOR

Programma van de Dag voor Statistiek en OR 2017
HEALTH CARE FOR THE FUTURE

Een bevende beurs

Mensen maken fouten, niet individuen

Wachtrijnetwerken en gelaagdheid

Loterijen onder vuur

Het ongepubliceerde boek van Lajos Takác

Kansrekening in de VS

Van Heraclitus tot Shannon;

De fluwelen revolutie van data in context en flux

Sexy Data Science;

Spiksplinternieuwe Sectie Data Science

M/V of V/M

IM Maarten van der Vlerk

Young Statisticians

Message from the President

Dear reader, 2017 has begun and soon, on the 23rd of March, we will celebrate the annual meeting of our society. In the preparation of the organisation of that meeting, the daily board of the VvS+OR took the initiative to meet with the boards of the individual VvS+OR sections as well as with societies dedicated to statistics and OR that are not part of the VvS+OR (e.g. Chemometric Society and Dutch Society for Ordination and Classification). The agenda for these meetings was as follows: 1. What does the section want to achieve, what are its main objectives?; 2. To which extent can the section achieve its objectives on its own, that is, independently of the VvS+OR?; 3. How can the VvS+OR help the section to achieve its objectives, does the VvS+OR have added value for the section?; 4. What should be the objectives of the VvS+OR in the near future? As can be read, the meetings were intended to clarify the role of the VvS+OR in relation to the sections and other societies.

An important conclusion of the meetings described above was that the annual meeting of the VvS+OR should, as much as possible, reflect joint interests and involve joint efforts. The annual meeting should also be a networking event where many statisticians and OR people from various backgrounds can meet and inform themselves about developments in their own field as well as in neighbouring fields, both for theory and practice.

Some further meetings and discussions between board and sections led to the topic of *Health Care for the Future* as a central theme for the annual meeting in 2017. This was considered to be an actual theme that elicits discussion at many levels in society, to name a few: politicians, policy makers, insurance companies, medical centres, citizens, etc. At this moment, the Dutch government spends 85 billion Euros per year on health care and it doesn't look as if this budget is open to strong increases while the demands do increase because of the aging of the population. Doctors did ask the politicians to make choices (*de Volkskrant*, 21 January 2017), for

example, do the Dutch want to improve general care for the elderly for a longer period of their life, or do they want to invest in the development of cures for relatively rare diseases?

Statistics + OR can contribute in important aspects to this debate on where to go with health care. New medical techniques require new statistical + OR methods for improving the medical practice or to make that practice more efficient. Quantitative methods are also important for the assessment of the quality of medical interventions and protocols; they help in evaluating and comparing scenarios. Furthermore, statistics + OR can develop indicators for monitoring the quality and efficiency of health care.

It is obvious that health care is a theme that touches on the interests of many people and any contribution to improving its quality and efficiency will attract the attention of a diverse audience. Therefore, health care is a theme by which the VvS+OR can show the Dutch audience that statistics + OR are highly relevant to science and society. The VvS+OR wants to increase the recognition and visibility of statistics + OR as a field and simultaneously of the involved scientists. As a consequence, additional funds may go to research for statistics + OR and more students may enroll in training and education programs, which in the end should lead to a higher level of statistics + OR in The Netherlands.

On the 23rd of March, thanks to the contributions of the sections of the VvS+OR, we can present an excellent program for our annual meeting. The important role of statistics + OR in the development and evaluation of health care will be demonstrated convincingly, thereby giving ample attention to scientific and policy aspects. Elsewhere in *STAtOR* you can find the speakers, titles and abstracts. I hope to welcome many of you on the 23rd of March.

FRED VAN EEUWIJK

HEALTH CARE FOR THE FUTURE

Thursday, March 23th 2017

Annual meeting of the Netherlands Society for Statistics and Operations Research (VvS+OR)

LOCATION
Jaarbeurs Utrecht
Jaarbeursplein 6, 3521 AL Utrecht (adjacent to Central Station)

PROGRAM

09:15 – 09:45	Registration and coffee tea
09:45 – 10:00	Opening
10:00 – 10:45	MAX WELLING (UvA)
10:45 – 11:15	HARWIN DE VRIES (Insead)
11:15 – 11:30	Break
11:30 – 12:00	ONNO VAN HILTEN (CBS)
12:00 – 12:30	Annual General Meeting (ALV, in Dutch)
12:30 – 13:30	Lunch break (at your own expense)
13:30 – 14:00	JOHAN LINSSEN (GGZ Momentum)
14:00 – 14:30	ELS GOETHGEBEUR (Ghent University)
14:30 – 15:15	Ceremony of the VAN ZWET AWARD
	Ceremony of the HEMELRIJK AWARD
	Final DATA SCIENCE HACKATHON
15:15 – 15:45	Break
15:45 – 16:30	MARK VAN DER LAAN (UC, Berkeley)
16:30	Snacks and drinks

1

SELF LEARNING ALGORITHMS IN HEALTH CARE

Max Welling

University of Amsterdam

ABSTRACT

After the sensational loss of the world champion GO against a computer algorithm, more and more people become aware of the fact that in many areas artificial intelligence could prove to be superior to human intelligence and intuition. Similarly in health care. After all, the human mind can process only a limited amount of information and can only detect simple patterns in data. Algorithms perform already better, for example, in analysing MRI images for predicting Alzheimer's disease, analysing histological images for predicting the gradation of a tumor, or analysing dermatological images for detecting melanomas. We are on the verge of a revolution where algorithms will enable the diagnosis and recommend the best treatments.

What does it take to make this revolution in health care to become reality? In this talk, I will highlight some developments in my discipline, machine learning, that are relevant for this process. For example, what is this technology that can analyse medical images so accurately? Can we train algorithms on patient data in such a way that they ensure privacy? Can we discover causal relationships from data without carrying out double-blind randomized studies? Can we automatically monitor patients and make recommendations 'on the fly'? What does it take to let physicians actually embrace this technology?

MAX WELLING has a research chair in Machine Learning at the University of Amsterdam and secondary appointments as full professor at the University of California Irvine and as a senior fellow at the Canadian Institute for Advanced Research (CIFAR). He is co-founder of Scyfer BV, a university spin-off in deep learning.

2

OPTIMIZING POPULATION SCREENING FOR INFECTIOUS DISEASES

Harwin de Vries

INSEAD School of Business, Fontainebleau

ABSTRACT

Population screening by mobile units is crucial to control several infectious diseases. We consider the following planning problem: given a set of populations at risk, the expected evolution of the epidemic in these populations, and a fixed number of mobile units, which populations should be screened when? We present descriptive models for the development of the burden of disease over time which take screening explicitly into account, use these to develop planning policies and solution methods, and numerically analyze them in the context of sleeping sickness control in D.R. Congo.

HARWIN DE VRIES studied econometrics at EUR. He received the Best Econometrics Thesis Award 2012. After earning his Ph.D. under Albert Wagelmans he now works for Technology and Operations Management, INSEAD School of Business, Fontainebleau, France.

3

HEALTH (CARE) STATISTICS IN THE NETHERLANDS

Onno van Hilten

Statistics Netherlands | CBS

ABSTRACT

Statistics Netherlands (SN) produces statistical information for the whole Dutch society, in particular for policy makers and researchers. In the field of health statistics, the classic examples are: causes of death statistics (since 1900), statistics on health, lifestyle and health consumption based on a national Health Interview Survey (since 1981), and statistics on health care expenditures (since 1972). In the last 15 years many new, very detailed statistics have been developed, based on 'statistical re-use' of administrative data which are generated by the primary processes in health care.

In the presentation examples will be given of the best long time series SN has on the one hand, and the new possibilities of all these recently developed statistics based on administrative data on the other hand. Furthermore, the question will be addressed to what extent all this statistical information enables researchers and policy makers to draw conclusions on accessibility, quality, effectiveness, efficiency and affordability of health care, and to compare health care in the Netherlands to other countries? Finally, the opportunities and challenges of the use of new big data sources for better health care statistics will be discussed.

ONNO VAN HILTEN graduated in mathematics in 1986, at the University of Groningen. He completed his PhD in mathematical economics in 1990, at the University of Maastricht. After working for 11 years at the Dutch National Energy Centre (ECN), he joined Statistics Netherlands (CBS) in 2001. Currently he is program manager 'health and health care statistics'.

4

OPERATION SUCCEEDED, PATIENT DIED

Johan Linssen

GGZ Momentum, Nijmegen

ABSTRACT

Momentum is a specialized psychiatric hospital group based in The Netherlands and South Africa. Momentum treats patients with psychiatric problems such as trauma and personality disorders. It also focuses on research & education, working closely together with Universities in Western Europe and Sub-Saharan Africa.

For a long time, researchers and mental healthcare specialists have upheld a hate-love relationship; 'protecting' patients from those *research nerds*. They also faced difficulties of coming up with classification of statistical methods in psychiatric care that are both simple and precise and that are able to address the complexity and unruly reality of mental health care problems. Regardless: protocols and manuals came forth from statistical research – incentivizing research on illnesses, rather than the patient itself – and thus created a larger distance to the patient's complaints and more importantly has not led to 'more health'. Every mental illness has a personal story and patient experience. Understanding that context is key to understanding how statistical research can improve the future of health care.

In essence, what will be discussed are not only the opportunities and challenges of using statistical research with innovating mental health care, but also the threats it poses to draw the wrong conclusions for those who seek help during their darkest hour.

JOHAN LINSSSEN is a Member of the Board at GGZ Momentum. Previous to that he was Country Director at consulting firm Berenschot in South Africa. He started his career at Capgemini and the European Commission. During his career he has consulted a large number of public and private organizations (such as Philips Healthcare and Departments of Health) on their organizational development. Johan has a background in economics. He is slightly claustrophobic and thus would hate to be a Pokémon.

5

COST EFFECTIVE MONITORING OF QUALITY OF CARE: THE GOLDMINE OF DISEASE REGISTERS AND HOW MORE CAN BE LESS

Els Goethgebeur

Ghent University

ABSTRACT

Public health infrastructure is under pressure. The gap between what is theoretically possible and practically affordable widens. Care centers receive populations of an increasingly advanced age, suffering from comorbidities. An evidence base for generalizable (average) effects of varying drug combinations requires causal inference and creative designs – including randomization within cohorts – but still leaves imprecise effect measures. Beyond this, there is a call for personalized medicine.

Informed by both general and local statistics from electronic health records, new directions for dynamic treatments are taken aiming at the right balance between quality of life, affordability and longevity. To evaluate varying strategies implemented in different care centers, outcomes that matter to patients should be monitored and point to center performance and room for improvement. Confounder adjustment and standardized outcomes per center will need the right balance of number and sophistication of patient specific covariates. Beyond standard statistical issues of bias and precision, one must account for the associated cost including registration fatigue with more measurement error and more missing data as a result. We discuss an approach for a cost-effective selection of covariates and the bias trade off balancing more confounders against additional measurement error and selectivity due to missing data.

ELS GOETGHEBEUR is professor of Statistics at Ghent University (B). After her graduate training in mathematics (KUL) and statistics (LUC) she held faculty positions at the London School of Hygiene and Tropical Medicine and Maastricht University. She taught at the Harvard School of Public Health and at Stanford University. Her research focuses on causal inference generally, on survival analysis and missing data problems, and more recently on big data problems in genetics.

6

TARGETED MACHINE LEARNING FOR CAUSAL INFERENCE: HARNESSING THE POWER OF BIG DATA TO IMPROVE HEALTH

Mark van der Laan

University of California

ABSTRACT

We review targeted minimum loss estimation (TMLE), which provides a general template for the construction of asymptotically efficient plug-in estimators of a target estimand for infinite dimensional models. TMLE involves maximizing a parametric likelihood along a so-called least favorable parametric model through an initial estimator (e.g., ensemble super-learner) of the relevant functional of the data distribution. The asymptotic normality and efficiency of the TMLE relies on the asymptotic negligibility of a second-order term. This typically requires the initial estimator to converge at a rate faster than $n^{-1/4}$. We propose a new estimator, the Highly Adaptive LASSO (HAL), of the data distribution and its functionals that converges at a sufficient rate regardless of the dimensionality of the data/model, under almost no additional regularity. This allows us to propose a general TMLE that is asymptotically efficient in great generality. We demonstrate the practical performance of HAL and its corresponding TMLE for the average causal effect. We also provide a review of some of the many applications of TMLE we have been involved in, involving the analysis of complex observational and experimental longitudinal studies, such as evaluation of impact of different treatment rules for controlling glucose level for diabetes patients.

MARK J. VAN DER LAAN is the Hsu/Peace Professor of Biostatistics at the University of California, Berkeley School of Public Health. He is the recipient of the 2005 COPSS Presidents' and Snedecor Awards, as well as the 2004 Spiegelman Award. His methodological research interests include censored data, causal inference, machine learning, multiple testing, semiparametric estimation theory, and targeted learning.

7

THE HACKATHON

Daniel Oberski

Utrecht University

ABSTRACT

A new activity in the VvS+OR is a hackathon. This hackathon has been inspired by the theme *Health care for the future*, with special thanks to prof. dr. Folkert Asselbergs, MD, from Utrecht University and University College London. Heart attacks are a major cause of mortality and we can therefore do a lot of good to society if we're able to prevent them. One of the main ways this is currently done is with medication such as statins; however, these pills, while preventing heart attacks, have the unfortunate side-effect of increasing the risk of diabetes. A highly important question is therefore 'can we prevent heart disease without increasing the risk of diabetes?'. To look at this, we will have at our disposal a dataset, cleverly constructed from publicly available sources by Asselbergs and his team, that combines negative health outcomes with genetic information. At the hackathon, teams will be invited to answer a few targeted questions formulated by the experts. These teams will be as mixed as possible, aiming to combine academic statisticians and computer scientists with data scientists working at companies and subject matter experts. During the awards ceremony, the jury will then discuss the work of the teams for all attendees of the Statistical Day. Thanks to the gracious support of the VvS+OR board, all hackathon participants will receive a free membership of the Society and the Section Data Science. (<http://sectiedatascience.nl/>)

DANIEL OBERSKI is associate professor of data science methodology at the department of Methodology & Statistics, Utrecht University. From 2006–2011 he worked at the ESADE business school and Pompeu Fabra University, Barcelona, on the European Social Survey. After obtaining his PhD from Tilburg University (2011) he became visiting assistant professor at the University of Maryland, US (2011), and subsequently postdoc and assistant professor at Tilburg University's (2012–2016). He is president of the VvS+OR Section on Data Science. Together with Katrijn Van Deun, Alessandro Di Bucchianico, Arend Oosterhoorn, Maarten Joosen and Erik-Jan van Kesteren, he hopes you will join this section and contribute to making it a succes for data lovers everywhere!

REGISTRATION (BEFORE MARCH 17th)

Members VvS+OR

The lectures and annual meeting are open to members of the Society, speakers and invited guests, but registration through the society's website is required: www.vvs-or.nl. Register before March 17th and bring the confirmation e-mail to show it at the registration desk at the entrance of the lecture room.

Non-members

Interested non-members can access the day and pay the amount of € 50 (bank account: NL42 INGB 0000 202091, BIC/ SWIFT: INGBNL2, VvS+OR 'registration annual meeting 2017') before March 17th and bring proof of payment (copy of bank statement) to show at the registration desk at the entrance of the lecture room. Alternatively, non-members can become an ordinary member for € 70 (go to www.vvs-or.nl and click 'Become a Member' in the top menu (young applicants will obtain a special rate), and have free access immediately after registration.

GENERAL INFORMATION

Language

The talks at the annual meeting will be in English, the Annual General Meeting (ALV) will be in Dutch.

Annual General Meeting (Algemene Ledenvergadering)

The Annual General meeting (ALV) is scheduled at the end of the morning. The relevant documents will be provided on the website two weeks before the meeting. You can also get them by e-mail if you send a request to info@vvs-or.nl.

Coffee, tea and drinks after the meeting

Coffee/tea during the breaks and drinks afterwards are offered by the Society.

Lunch

Lunch is at your own expense. The restaurant of the Jaarbeurs is usually quite busy, but within walking distance of just a few minutes (e.g. in the direction of Central Station) you will find many alternatives.

ORGANIZING COMMITTEE

The annual meeting is organized by the board of the VvS+OR. For questions, contact the administration by e-mail at info@vvs-or.nl.

STAtOR is een uitgave van de Vereniging voor Statistiek en Operationele Research (VvS+OR). STAtOR wil leden, bedrijven en overige geïnteresseerden op de hoogte houden van ontwikkelingen en nieuws over toepassingen van statistiek en operationele research. Verschijnt 4 keer per jaar.

Redactie

Joaquim Gromicho (hoofdredacteur), Annelieke Baller, Ana Isabel Barros, Joep Burger, Kristiaan Glorie, Caroline Jagtenberg, Guus Luijben (eindredacteur), Richard Starmans, Gerrit Stermerdink (eindredacteur) en Vanessa Torres van Grinsven. Vaste medewerkers: Johan van Leeuwen, Fred Steutel en Henk Tijms.

Kopij en reacties richten aan

Prof. dr. J.A.S. Gromicho (hoofdredacteur), Faculteit der Economische Wetenschappen en Bedrijfskunde, afdeling Econometrie, Vrije Universiteit, De Boelelaan 1105, 1081 HV Amsterdam, telefoon 020-5986010, mobiel 06-55886747, <j.a.dossantos.gromicho@vu.nl>.

Bestuur van de VvS+OR

Voorzitter: prof. dr. Fred van Eeuwijk <president@vvs-or.nl>
Secretaris: dr. Laurence Frank <bestuur@vvs-or.nl>
Penningmeester: dr. Ad Ridder <penningmeester@vvs-or.nl>
Overige bestuursleden: prof. dr. Eric Cator (SMS), prof. dr. Jeanine Houwing-Duistermaat (BMS), Maarten Kampert MSc., prof. dr. Albert Wagelmans (NGB), dr. Michel van de Velden (ECS), dr. Jelte Wicherts (SWS), Kees Mulder (Young Statisticians)

Leden- en abonnementenadministratie van de VvS+OR

VvS+OR, Postbus 1058, 3860 BB Nijkerk, telefoon 033 - 2473408, e-mail <admin@vvs-or.nl>.
Raadpleeg onze website over hoe u lid kunt worden van de VvS+OR of een abonnement kunt nemen op STAtOR.

VvS+OR-website

www.vvs-or.nl

Advertentieacquisitie

M. van Hootegem <hootegem@xs4all.nl>
STAtOR verschijnt in maart, juli, oktober en december.

Ontwerp en opmaak

Pharos, Nijmegen

Uitgever

© Vereniging voor Statistiek en Operationele Research
ISSN 1567-3383

Nieuw verbeterd!

Reclame doet ons geloven dat wij voortdurend behoefte hebben aan iets nieuws. En dat nieuwe moet vooral verbeterd zijn! Of het nu een nieuwe versie van tandpasta betreft of een elektrische fiets, de reclame suggereert dat we sukkelers zijn omdat we al die tijd de oude niet verbeterde versie te hebben geaccepteerd.

Is dat gedrag inmiddels ook waarneembaar in ons vakgebied? Gelukkig niet. Weliswaar schrijft Daniel Oberski zeer enthousiast over een 'Spiksplinternieuwe Sectie Data Science' binnen onze vereniging, maar hij benadrukt vooral dat het hier gaat om het combineren van kennis, inzicht en inspiratie uit verschillende takken van ons mooie vak. En dat combineren is niet 'verbeterd', maar een gebruikelijke werkwijze in de wetenschap.

Ook het – lange maar zeer leeswaardige – artikel van Richard Starmans noemt regelmatig nieuwe zaken, maar geeft toch vooral een overzicht van het tot stand komen van nieuwe inzichten op basis van kruisbestuiving vanuit meerdere disciplines.

Dit nummer van STAtOR is vooral belangrijk als aankondiging van de Dag voor Statistiek & OR 2017. Die Dag zou eveneens als 'Nieuw, verbeterd!' kunnen worden aangeprezen maar ook hier mist u de typische reclamekreten. Het bijzonder interessante programma biedt een breed scala van presentaties die allemaal gefocust zijn op het toepassingsgebied Zorg. Zij geven goed voorbeeld van hoe ontwikkelingen in ons vakgebied kunnen bijdragen aan de zeer diverse problemen in de zorg.

Die aandacht voor toepassingen in de zorg is overigens niet nieuw, en ook niet verbeterd! Als u terugbladert in de oude jaargangen STAtOR zult u veel artikelen hierover tegenkomen.

Misschien is dat wel de juiste aanpak: niet steeds roepen dat iets 'Nieuw, verbeterd!' is, maar wél regelmatig even stilstaan om vast te stellen of we iets van elkaar en van andere vakgebieden kunnen leren. Dan komt de vernieuwing vanzelf, al zullen we die soms pas in een terugblik herkennen.

De redactie wenst u veel ouderwets leesplezier!

DE REDACTIE



INHOUD

Dag voor Statistiek en OR 2017

2 Message from the President | FRED VAN EEUWIJK

3 Health Care for the future
Programma Dag voor Statistiek en OR 2017

8 Redactioneel

10 Een bevende beurs | FRANCINE GRESNIGT

14 Mensen maken fouten, niet individuen | CHRIS HARTGERINK

17 Wachtrijnetwerken en gelaagdheid | JAN-PIETER DORSMAN

20 Loterijen onder vuur – column | HENK TIJMS

22 Van Heraclitus tot Shannon: de fluwelen revolutie van data in context en flux | RICHARD STARMANS

31 Agenda

32 Het ongepubliceerde boek van Lajos Takács | ONNO BOXMA

33 Kansrekening in de VS – column | FRED STEUTEL

34 Sexy Data Science; spiksplinternieuwe Sectie Data Science | DANIEL OBERSKI

37 M/V of V/M – column | GERRIT STERMERDINK

38 IM Maarten van der Vlerk | WIM KLEIN HANEVELD & WARD ROMEIJNDERS

39 Young Statisticians



Reddingswerkers aan de slag in Pescara del Tronto in Italië na de aardbeving op 24 augustus 2016

FRANCINE GRESNIGT

De volgende twee nieuwsberichten zullen u wellicht niet ontgaan zijn.

De Amerikaanse beursgraadmeters verloren donderdag voor de zevende handelsdag op rij veel terrein, door aanhoudende zorgen over de gevolgen van de kredietcrisis. Beleggers vrezen dat de woensdag aangekondigde gecoördineerde rentverlagingen door centrale banken wereldwijd een recessie niet kunnen tegenhouden. De Dow-Jones index van dertig toonaangevende fondsen sloot met een verlies van 678,91 punten (7,3 procent) op een stand van 8579,19 punten, het laagste niveau in meer dan vijf jaar. De S&P 500-index gaf met een verlies 75,02 punten, ofwel 7,6 procent, op 909,92 punten nog wat meer terrein prijs. De technologie-index Nasdaq moest de winst van eerder op de dag inleveren en verloor 95,21 punten (5,5 procent) op 1645,12 punten. 'Er is nog een belangrijke nervositeit over de duur van de kredietcrisis en de problemen die nog voor de boeg liggen', aldus een handelaar in New

York. 'Beleggers zijn bang voor wat ze niet weten en die angst zet ze aan tot de verkoop van hun aandelen. Vragen stellen is voor later.' [...] (*de Volkskrant*, 9 oktober 2008)

En acht jaar later.

Midden-Italië is in de nacht van woensdag op donderdag opnieuw getroffen door een aardbeving. De regio die recentelijk al door meerdere bevingen werd getroffen, kreeg dit keer een aardbeving met een kracht van 4,8 te verduren, zo meldt het Italiaanse persbureau ANSA. [...] Zondag vond er een aardbeving met een kracht van 6,5 plaats in de buurt van Norcia, de zwaarste aardbeving in het land in 35 jaar tijd. Enkele dagen ervoor waren er ook al bevingen van 5,5 en 6,1 geweest. Sindsdien zijn er meer dan duizend kleine naschokken gemeten.' (*de Volkskrant*, 3 november 2016)

Ogenschijnlijk lijken deze berichten niets met elkaar te maken te hebben. Wanneer we de berichten echter beter

EEN BEVENDE BEURS

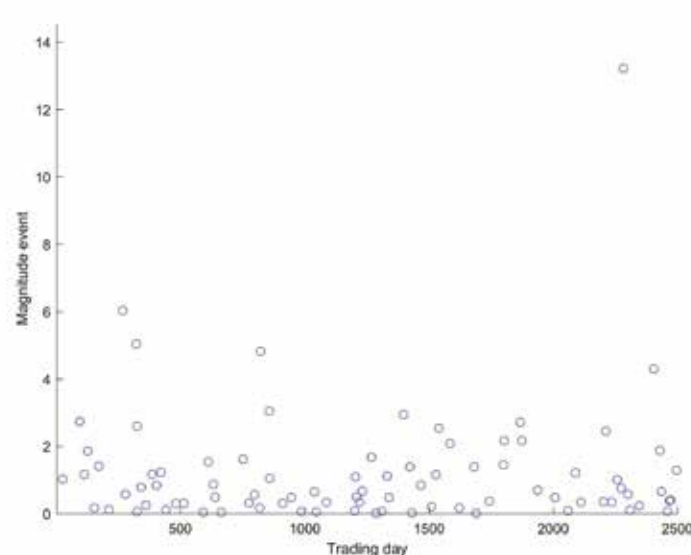
Het gedrag van aandelen-rendementen rondom een financiële crash lijkt op het gedrag van de seismische activiteit rondom aardbevingen. Beiden zorgen ervoor dat schokken elkaar uitlokken en clusteren in zowel tijd als ruimte/de cross-sectie waarbij de omvang van een schok invloed heeft op het aanstekende effect van de schok. Op basis van deze gelijkenissen ontwikkelen we een model voor het voorspellen van een beurscrash. Wanneer we dit model gebruiken in een waarschuwingssysteem voor crashdagen en testen op de S&P 500 index gedurende de recente financiële crisis vinden we dat het model in staat is om de kans op crashes op de middellange termijn te voorspellen.

bestuderen, kunnen we enige gelijkenis ontdekken: zowel op de Amerikaanse beurzen als op de randen van de tektonische platen die in Italië bij elkaar komen, vinden schokken plaats. Deze schokken vormen een reeks en vinden plaats in een bepaalde periode, voorafgaand en gevolgd door perioden waarin weinig tot geen schokken plaatsvinden. De schokken worden niet alleen gevoeld op één beurs of één plaats, maar op meerdere beurzen en plaatsen. Deze beurzen of plaatsen staan met elkaar in verbinding omdat op beiden dezelfde bedrijven en/of beleggers actief zijn, dan wel beiden geografisch gezien dichtbij elkaar liggen. De periode waarin veel schokken plaatsvinden kan worden aangemerkt als 'een gespannen periode', waarin er spanningen zijn ofwel bij aandeelhouders dan wel in de ondergrond waar aardplaten bij elkaar komen. Dit brengt ons bij de vraag: Gaat achter de bevingen op de beurs hetzelfde mechanisme schuil als achter aardbevingen?

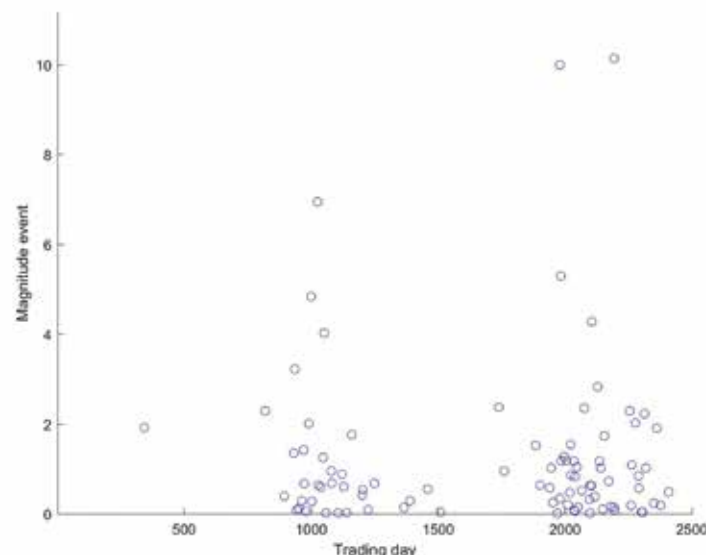
Als er één schaap over de dam is, dan...

De identificatie en voorspelling van crashes is erg belangrijk voor handelaren, toezichhouders van financiële markten en risicomanagement, omdat een serie van omvangrijke negatieve prijsveranderingen gedurende een korte periode ernstige gevolgen kan hebben. Tijdens de

kredietcrisis leden financiële indices talrijke handelsdagen enorme verliezen, wat in 2008 leidde tot het grootste jaarlijkse procentuele verlies voor zowel de Dow-Jones index, de S&P 500 index en de Nasdaq sinds 1931. De crisis demonstreerde de samenhang tussen de indices. Op 29 september, 15 oktober en 1 december 2008 leden alle drie de indices procentuele verliezen behorende tot de top 20 meest negatieve rendementen op de indices. Om een model te ontwikkelen voor beurscrashes focussen we ons eerst op de mogelijke oorzaken van zulke crashes. Sornette (2003) vat samen dat algorithmisch handelen, de toegenomen handel in derivaten, illiquiditeit, handels- en budgettekorten en overwaardering opeenvolgende negatieve koersbewegingen kunnen provoceren. Belangrijker, Sornette (2003) wijst erop dat speculatieve zeepbellen, welke leiden tot crashes, het gevolg zijn van het kuddegedrag van beleggers. Dit kuddegedrag zorgt ervoor dat crashes zichzelf verergeren. Dus, terwijl het uiteenspatten van de speculatieve zeepbellen getriggerd kan worden door exogene omstandigheden (nieuwsberichten bijvoorbeeld), groeit instabiliteit op beurzen intrinsiek. Een model voor beurscrashes moet daarom in staat zijn om deze zelf-excitatie te vatten. Een dergelijke zelf-excitatie kan ook worden opgemerkt in het seismische gedrag rond series aardbevingen, waarbij een aardbeving naschokken genereert, die op hun beurt weer naschokken genereren.



Figuur 1. Geen zelf-excitatie



Figuur 2. Zelf-excitatie

Een model voor geclusterde schokken

In de geofysische literatuur wordt al geruime tijd onderzocht of statistische modellen de seismische activiteit in een regio kunnen omschrijven. Een van de meest populaire modellen vanwege zijn eenvoud, interpreteerbaarheid maar bovenal goede fit, is het Epidemic Type Afterschock Sequence model (Ogata, 1988). In het ETAS-model is de plaatselijke frequentie waarmee aardbevingen optreden een functie van de voorafgaande seismische activiteit. Deze functie heeft een specifieke vorm, welke gebaseerd is op empirische wetten in tijd en ruimte zoals een gemodificeerde versie van de Omori-wet, de Gutenberg-Richter-wet en de Utsu-Seki-wet. De Omori-wet beschrijft de frequentie waarmee naschokken plaatsvinden gegeven de tijd na de hoofdschok. De Gutenberg-Richter-wet relateert de grootte van een aardbeving aan de hoeveelheid naschokken. De Utsu-Seki wet geeft aan hoe het gebied dat gevoelig is voor naschokken zich verhoudt tot de grootte van een eerste schok. Volgens deze wetten neemt de hoeveelheid naschokken in een gebied, veroorzaakt door een aardbeving, af naarmate er tijd verstrijkt, de aardbeving kleiner is en het gebied verder ligt van het epicenter van de aardbeving. Echter, naschokken kunnen weer nieuwe naschokken veroorzaken. Bovendien hoeft de eerste schok niet de grootste schok te zijn. De functie, welke de seismische activiteit op een bepaald tijdstip in een bepaald gebied weergeeft, is een optelsom van de invloeden op de activiteit van de alle aardbevin-

gen tot op dat moment. Deze functie kan gebruikt worden om de kans op naschokken te schatten. Werkelijke tijdstippen, plaatsen en aantallen van naschokken zijn stochastisch en dus niet vooraf te bepalen.

Schok op schok

Terug naar de beurs. In het ETAS-model neemt de kans op schokken toe wanneer er een schok plaatsvindt, waarna deze afneemt als functie van de tijd na de schok. Zelfexcitatie kan dus door het ETAS-model gevat worden. De specifieke vorm van de functie, welke de invloed van verschillende voorgaande crashes op de kans op een extreme prijsdaling beschrijft, aangepast worden. Op die manier kan het model de geobserveerde relaties tussen de hoeveelheid door een crash geïnduceerde prijsdalingen en de grootte van de crash, en het vermogen van de crash om een crash uit te lokken in een andere index (cross-excitatie, Gresnigt et al., 2017a en 2017b) vatten.

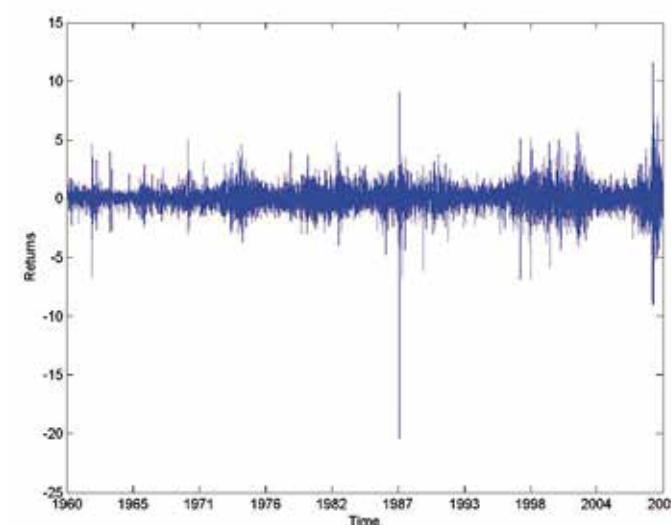
Het ETAS-model lijkt dus niet alleen bruikbaar te zijn voor aardbevingen. Het model is eerder onder andere toepast op data van misdaden (Mohler et al., 2011) en rode bananenplanten (Balderama et al., 2012). Kenmerkend voor data welke geschikt zijn om gemodelleerd te worden door het ETAS-model is het clusteren van extremen in tijd en ruimte/de cross-sectie. Dit is te zien in de figuren 1 en 2.

Een waarschuwingssysteem voor crashes

Gebruik makend van het ETAS-model kunnen we de kans op een crash in een index binnen een bepaald tijdsinterval voorspellen (Gresnigt et al., 2015). Op deze manier stelt het ETAS-model ons in staat te waarschuwen voor een aankomende crash (aardbeving) in de index. Om een waarschuwingssysteem voor toekomstige beurscrashes te ontwikkelen, is het nodig dat we eerst de parameters van het ETAS-model bepalen. Deze parameters definiëren vervolgens de kans op crashes gegeven de crash historie van de beurs (de tijdstippen en de omvang van eerdere crashes).

Wanneer we het ETAS-model toepassen op de 5% meest negatieve rendementen op de S&P 500 index over de periode van 2 januari 1957 tot 1 september 2008, bevestigen de schattingsresultaten dat, net als aardbevingen, crashes zichzelf aanwakkeren. De snelheid waarmee de kans op het uitlokken van nieuwe crashes afneemt als functie van tijd sinds een crash, ligt echter hoger dan bij aardbevingen. Gelijkelijk aan aardbevingen is de hoeveelheid crashes die een crash veroorzaakt groter wanneer de omvang van de crash groter is. Een andere bevinding is dat, gemiddeld genomen, de omvang van crashes groter is gedurende een periode waarin meer en/of grotere crashes plaatsvinden dan in een rustiger periode. Figuur 3 bevat een grafiek van de gemodelleerde crash intensiteit. De pieken gedurende de krediet crisis en rondom Black Monday bevestigen onder andere dat het ETAS-model de tijdsvariatie in de frequentie waarmee crashes plaatsvinden goed kan vatten.

Met behulp van de geschatte parameters ontwikkelen we een waarschuwingssysteem voor crashes in de index gebaseerd op de kans dat een crash plaatsvindt binnen een handelsweek. We testen het waarschuwingssysteem



Figuur 3. Crash-intensiteit

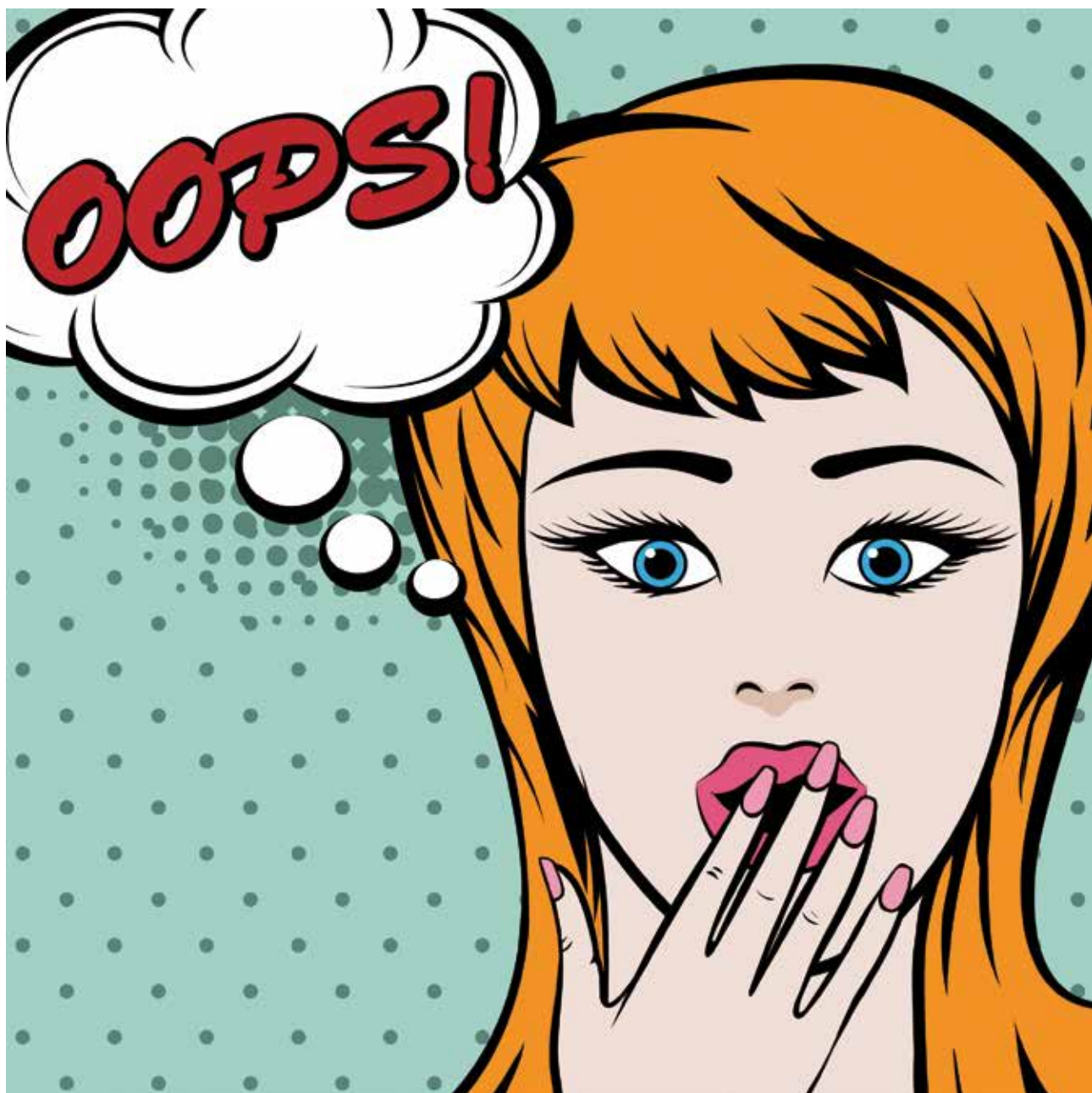


van 1 september 2008 tot 1 januari 2013, waarbij we vinden dat de frequentie van correcte voorspellingen van crashes hoger ligt dan de frequentie van valse voorspellingen. Hieruit concluderen we dat het modelraamwerk dat het zichzelf versterkende gedrag van rendementen rondom financiële crashes uitbuit, in staat is om de kans op crashes op de middellange termijn te voorspellen.

LITERATUUR

- Balderama, E., Schoenberg, F.P., Murray, E., & Rundel, P.W. (2012). Application of branching models in the study of invasive species. *Journal of the American Statistical Association*, 107(498), 467-476.
- Gresnigt, F., Kole, E., & Franses, P.H. (2015). Interpreting financial market crashes as earthquakes: A new Early Warning System for medium term crashes. *Journal of Banking & Finance*, 56(1), 123-139.
- Gresnigt, F., Kole, E., & Franses, P.H. (2017a). Specification Testing in Hawkes models. *Journal of Financial Econometrics*, forthcoming.
- Gresnigt, F., Kole, E., & Franses, P.H. (2017b). Exploiting Spillovers to forecast crashes. *Journal of Forecasting*, forthcoming.
- Mohler, G.O., Short, M.B., Brantingham, P.J., Schoenberg, F.P., & Tita, G.E. (2012). Self-exciting point process modeling of crime. *Journal of the American Statistical Association*, 106(493), 100-108.
- Ogata, Y. (1988). Statistical models for earthquake occurrences and residual analysis for point processes. *Journal of the American Statistical Association*, 83(401), 9-27.
- Sornette, D. (2003). Critical market crashes. *Physics Reports*, 378(1), 1-98.

FRANCINE GRESNIGT is promovendus aan het Econometrisch Instituut van de Erasmus Universiteit Rotterdam. Tevens is zij lid van het Tinbergen Instituut. E-mail: gresnigt@ese.eur.nl



Illustratie: neyro2008

Mensen maken fouten, niet individuen

Ondanks dat we als individu er veel voor doen om geen fouten te maken in ons werk, sluipen ze er toch regelmatig in. Zo heb ik bijvoorbeeld geprobeerd zo min mogelijk spel- en taalfouten te maken in dit artikel. Desondanks zal mij er een aantal ontglipt zijn, deels omdat het inefficiënt is perfectionistisch daarin te zijn en deels omdat de onopgemerkte fouten waarschijnlijk minimaal zijn. Daarom heb ik een spellingchecker gebruikt die nog een aantal fouten aankaarte op efficiënte wijze. Met fouten maken is niks inherent mis; mensen zijn als soort nu eenmaal feilbaar. De belangrijk vraag is: erkennen we dat we fouten maken en stellen we ons daar op in?

CHRIS HARTGERINK

Binnen het huidige hypercompetitieve wetenschapsbedrijf ontstaat regelmatig een wetenschapper die vooral bezig is met het imago en de realiteit als zekerder schetst dan dat we weten dat deze is. Een fout maken wordt dan ook als een aberratie gezien van het geveinsde zekere beeld, waarna dit als een gevaar voor het imago en reputatie ervaren kan worden. Hierdoor komt vervolgens ook nog eens mogelijk het competitieve voordeel in het geding. Het wordt op deze manier erg lonend om fouten weg te moffelen en mogelijke fouten niet te delen. Hierdoor ontstaat er ook meteen een taboe op fouten maken voor jonge onderzoekers die nog zoekende zijn binnen het systeem; bijna niemand heeft het immers over fouten maken. Waar is de kritische blik naar binnen gebleven, vraag ik mij dan af? Uiteraard, het is niet altijd zoals ik hier schets, maar ik denk dat menig lezer zich eerder zou schamen voor een fout dan het als een ontdekking te zien die begrepen dient te worden.

'Als fouten de conclusies niet aantasten, ga ik ze niet corrigeren want ze doen er niet toe' wordt dan regelmatig ge(m)opperd. Echter, als ik hierover nadenk klinkt een enorme zelfingenomenheid door: de wetenschapper weet wel hoe het moet en de rest moet daarop vertrouwen. Sterker nog, een dergelijke redenering maakt het een voldongen feit; de beoordeling van de fout door de wetenschapper moet maar geaccepteerd worden. Helaas is dit echter een façade die hopelijk een keer ten einde komt — wanneer een substantieve conclusie niet wijzigt zullen de schattingen, precisie, etc. dat als nog wel doen. En laat het natuurlijk weer net zo zijn dat juist op deze resultaten wordt voortgebouwd in verder onderzoek of waarop meta-analyses van een fenomeen gebaseerd worden. Op termijn veroorzaakt dit een immense afwijking van de werkelijkheid, door allerlei kleine stapjes in de verkeerde richting.

Fouten toegeven komt dan ook ten goede aan de wetenschap in zijn geheel, fouten wegmoffelen komt ten goede aan het individu en zijn imago. Daarnaast is wegmoffelen voordelig omdat het simpelweg te makkelijk is om te doen en het risico dat iemand er achter komt laag is. Artikelen worden opgesloten achter peperdure abonnementen, materialen worden niet (of niet goed) gedocumenteerd of gearhiveerd en data worden bewaakt als eigendom terwijl zij dat niet zijn. Hierdoor wordt fouten ontdekken moeilijker. Als we openheid eisen over zaken als samenleving en beroepsgroep, zullen fouten niet weg te moffelen zijn en moeten we ze omarmen. Dit zal het eigenbelang gelijk laten lopen met het systeembelang: fouten toegeven komt ten goede van de wetenschap en het individu. Vervolgens zal het imago bepaald worden door hoe men met fouten omgaat, iets waar iedereen zelf over kan beslissen.

Daarnaast wordt fouten maken een simpele *sunk cost fallacy* door imago-denken; erkennen is afhankelijk van het imago en de competitie, niet van de fout zelf. Ik verwacht dan ook dat jonge onderzoekers eerder geneigd zijn fouten toe te geven. Maar laten we onszelf niet voor de gek houden: we trappen allemaal in de *sunk cost fallacy*, zelfs wanneer we ervan weten. Alleen wanneer we waarborgen introduceren, zoals transparantie, kunnen we onze eigen feilbaarheid voor zijn.

Foutief rapporteren

Fouten kunnen zich manifesteren in veel verschillende hoedanigheden; één specifieke variant is het simpelweg foutief rapporteren van statistische resultaten. Mijn collega Michèle Nuijten heeft hiervoor binnen haar PhD-project software ontwikkeld (samen met Sacha

Epskamp) die artikelen doorneemt en zoekt naar mogelijke rapportagefouten. Bijvoorbeeld: een onderzoeker rapporteert $t(69) = 2,34, p < 0,01$. Deze software herkent het resultaat, verzamelt de verschillende onderdelen van het gerapporteerde resultaat en herberekent op basis van de vrijheidsgraden en toetswaarde de p-waarde. Wanneer de gerapporteerde en herberekende p-waarden overeenkomen, is er niks aan de hand; wanneer er een discrepantie is wordt het als een mogelijke inconsistentie aangegeven en kan de onderzoeker de fout corrigeren. Wanneer statistische significantie wijzigt ($\alpha = 0,05$), is er mogelijkwerwijs een statistische beslissingsfout gemaakt. In het voorbeeld is de p-waarde incorrect, maar de statistische conclusie op basis van het resultaat is correct; de p-waarde zou hier 0,022 moeten zijn. Simpelweg in de verkeerde kolom van de output kijken kan al genoeg zijn — een fout zit immers in een klein hoekje. Een dergelijk programma dient als waarborg voor een specifieke soort fout die onderzoekers kunnen maken.

Deze *statistics checker* – of kortweg *statcheck* – is open source software en kan door iedereen hergebruikt worden. Het heeft ons als onderzoeksmiddel al een aantal interessante bevindingen opgeleverd, die ons meer leren over de mens als wetenschapper. Zo blijkt dat het aantal fouten stabiel is over de jaren, waar ongeveer 1 op 10 resultaten een fout bevat; dit alles ondanks dat we zouden verwachten dat computers de mens steeds preciezer kunnen maken op dit vlak. Ook leren we hieruit dat meer resultaten rapporteren leidt tot meer fouten in absolute zin, maar niet in relatieve zin. Daarnaast heeft het blijkbaar weinig zin om de verantwoordelijkheid voor de resultaten te delen; er worden niet systematisch minder fouten gemaakt wanneer meer auteurs deze verantwoordelijkheid dragen. Deze systematiek lijkt te duiden dat het menselijke proces feilbaar is en niet het specifieke individu.

Wat kunnen we dan doen om fouten te voorkomen? *statcheck* is niet alleen een onderzoekstool; het is ook een tool die elke onderzoeker kan gebruiken om een eigen manuscript door te lichten op foutjes (via *statcheck.io* of in R). Op het moment opereert *statcheck* het beste binnen de psychologie voor t-, F-, r-, z- en Chi-kwadraat-toetsen. Wanneer deze toetsen in één van uw manuscripten zit, raad ik u dan ook aan het een keer uit te proberen. Het is zo gepiept. En laten we eerlijk zijn, het overtypen van statistische resultaten naar een manuscript is vervelend en nogal repetitief, dus we kunnen wel een hulpmiddel gebruiken wanneer we met net iets teveel afleiding cijfers hebben ingeklopt.

Menselijk proces

Genezen is echter nog beter dan voorkomen; *statcheck* voorkomt foutjes, maar geneest hun oorsprong, het menselijke proces, niet. Waarom zouden we nog handmatig al deze resultaten inkloppen? Het is immers een systematisch taakje en computers zijn fantastisch met systematische taken. Om deze reden zijn hulpmiddelen zoals RMarkdown geïntroduceerd: geef de analysecode aan die het resultaat produceert en het wordt dynamisch in je manuscript ingevoegd. Zo rapporteer je niet langer 'F(1,12) = 2,54, p = 0,137', maar 'F(res\$df1, res\$df2) = res\$f_value, p = res\$p_value'. Wanneer je je wordfile dan aanmaakt, worden automatisch de cijfers ingevoegd om het correcte resultaat te presenteren. De computer geneest niet alleen de oorsprong van de rapportagefouten, maar maakt ons leven zelfs makkelijker. Het heeft bij mij al menige fout voorkomen, maar het is geen absoluut middel: wanneer je de verkeerde informatie erin duwt, komt er alsnog verkeerde informatie uit.

50.000 auteurs

Gegeven de systematiek van fouten en dat ze soms onopgemerkt lijken te gaan, ondernam ik onlangs een project waarin ik 50.000 auteurs op de hoogte wilde brengen van mogelijke fouten in hun artikelen. Samen met de post-publicatie peer review website PubPeer plaatste ik korte, automatisch gegenereerde rapporten voor de artikelen die ik met *statcheck* doorgenomen had. Hiermee wilde ik auteurs attenderen op mogelijke foutjes en ook deze foutjes bespreekbaar maken. Ik had echter niet voorzien dat er zulk een grote discussie zou ontstaan over, eerlijk gezegd, kleine fouten. Mensen dachten dat ik ze beschuldigde van intentioneel misrapporteren of waren bang voor reputatieschade. Er was ook waardering, waar sommigen het zagen als een dienst. Eén ding dat het zeker heeft gedaan is de diversiteit van opinies over fouten naar voren brengen en ervoor zorgen dat het een keer toont hoe we als wetenschappers met fouten omgaan. Fouten maken is menselijk; hoe je ermee omgaat getuigt van hoe je je daar als mens overheen probeert te zetten.

CHRIS H. J. HARTGERINK is PhD-kandidaat aan de Tilburg University. Hij specialiseert zich in het detecteren van mogelijke datafabricatie en hoe menselijke factoren invloed hebben op het onderzoeksproces. E-mail: C.H.J.Hartgerink@uvt.nl
De tekst van dit artikel is beschikbaar onder de licentie CC0 1.0 Rights Waiver en is vrij herbruikbaar



WACHTRIJNETWERKEN EN GELAAGDHEID

JAN-PIETER DORSMAN

Niet-gelaagde netwerken van wachtrijen vormen een belangrijk onderwerp in de literatuur van de toegepaste kansrekening en de stochastische operations research dat zeer uitgebreid bestudeerd is. Dit is mede te danken aan de relevantie van deze modellen in tal van productie-, computer- en communicatiesystemen. In deze toepassingen is tegenwoordig echter steeds vaker sprake van een gelaagde structuur die een grote impact heeft op de prestatie van deze systemen. Onderzoek naar zogenaamde gelaagde wachtrijnetwerken wordt daarmee van groot belang.

Klanten en bedienden

In het midden van de twintigste eeuw introduceerde J.R. Jackson het concept van netwerken van wachtrijen (Jack-

son, 1957). In deze netwerken wachten klanten in wachtrijen totdat ze geholpen worden door een bediende. Nadat een klant bediend is, verlaat deze het systeem of sluit achter aan bij een andere wachtrij om daar geholpen te worden door de bij die rij behorende bediende. Op die manier ontstaat een netwerk van wachtrijen, dat onder bepaalde voorwaarden succesvol geanalyseerd kan worden. Zo wist Jackson een uitdrukking af te leiden voor de (gezamenlijke) kansverdeling van de hoeveelheden wachtende klanten in elke rij.

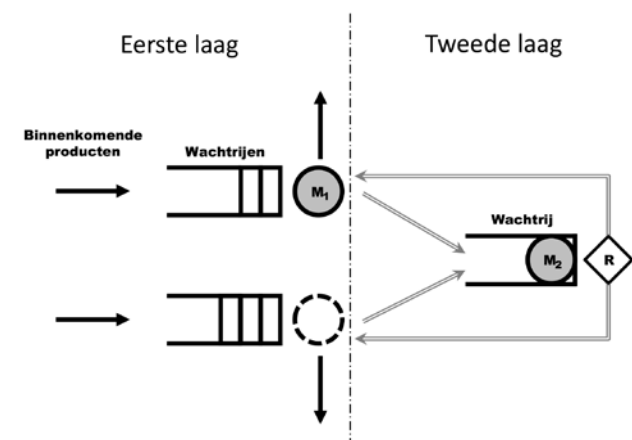
Sinds dit werk van Jackson is veel aandacht besteed aan wachtrijnetwerken vanuit verschillende invalshoeken. Vanuit theoretisch oogpunt is veel werk verricht aan het verzwakken van de modelaannames die Jackson deed voor zijn analyse. Daarnaast is met name vanuit de logistiek en de informatica uitgebreid onderzoek gedaan

naar de implicaties van deze resultaten bij het ontwerp of de optimalisatie van processen in de praktijk. Denk bijvoorbeeld aan het ontwerp van een productiesysteem in een fabriek, waarbij producten sequentieel door verschillende machines bewerkt moeten worden en tussentijds bij elke machine op hun beurt moeten wachten, of aan een computersysteem waarbij verschillende taken achtereenvolgens door een of meerdere processoren moeten worden uitgevoerd.

Vrijwel alle studies op dit gebied hebben gemeen dat er een strikt onderscheid wordt gemaakt tussen klanten en bedienden. Zo worden de producten en de taken in bovenstaande voorbeelden altijd aangeduid als klanten, aangezien dit entiteiten zijn die bediening behoeven van de machines en processoren. De machines en processoren worden dan ook eenduidig als bedienden aange-merkt.

Gelaagde structuur

In veel hedendaagse toepassingen blijkt de strikte scheiding tussen klanten en bedienden in wachtrijnetwerken een steeds moeilijker te rechtvaardigen aanname. Zo kan in het bovengenoemde productiesysteem een machine, ondanks zijn gedoodverfde rol als bediende, kapotgaan, waarna deze bediening behoeft door een reparateur. In die omstandigheid neemt de machine juist de rol aan van een klant met de reparateur als bediende. Dit wachtrijnetwerk krijgt op deze wijze een gelaagd karakter (figuur 1). Deze figuur laat twee machines zien. Machine M_1 vervult in de eerste laag van het systeem zijn



Figuur 1. Voorbeeld van een gelaagd wachtrijnetwerk, waarin machine M_1 bediende van producten is en machine M_2 klant is geworden van reparateur R.

bedieningsrol door wachtende producten te bewerken. Machine M_2 is echter onlangs kapot gegaan, waardoor deze zelf als klant heeft plaatsgenomen in een wachtrij in de tweede laag van het systeem om gerepareerd te worden door de reparateur R.

Overigens is dit slechts een van de vele mogelijke voorbeelden van gelaagde wachtrijnetwerken, die in de praktijk voor het oprapen liggen. Zo wordt een ander tot de verbeelding sprekend voorbeeld gevormd door peer-to-peernetwerken, waarbij gebruikers in hun rol als klant diensten genieten van anderen, maar tegelijkertijd een rol als bediende invullen door dezelfde (of vergelijkbare) diensten aan medegebruikers van het netwerk aan te bieden. Men kan ook denken aan een groot callcenter, waarbij een telefonist de behoefte kan hebben om een senior medewerker te raadplegen. In zo'n geval ruilt de telefonist tijdelijk zijn bedienderol in voor die van een klantrol. Dergelijke situaties treden op bij tal van dienstverlenende organisaties. Daarnaast zijn gelaagde structuren alomtegenwoordig in allerlei soorten computersystemen. Dezelfde gelaagdheid komt ook voor in minder verwachte hoeken. Zo gebruikt een overslagbedrijf in de haven automatisch gestuurde voertuigen (AGV's) om containers te vervoeren tussen een opslagplaats en de schepen. Hierbij is een kadekraan nodig om een container van de AGV naar het schip over te laden, of andersom. In zo'n geval is de AGV een bediende van de containers, maar gedraagt zich als klant bij de kadekraan.

Interacties tussen lagen

De wiskundige analyse van gelaagde wachtrijnetwerken blijkt een lastige aangelegenheid. Dit is vooral een gevolg van het feit dat de lagen in deze netwerken een intensieve wisselwerking hebben, zoals blijkt in het eerdergenoemde voorbeeld van het productiesysteem. In de eerste laag van het systeem is nog vrij goed te kwantificeren wat de invloed is die de machine heeft op zaken als de totale verblijftijd van een product (ofwel de totale tijd dat een product in een wachtrij/buffer zit en vervolgens in bewerking is). De machine heeft immers een directe rol als bediende van het product. Als de machine bijvoorbeeld twee keer zo snel zou werken, zou deze tijd ongeveer halveren. Echter, de vraag wat er zou gebeuren met de verblijftijd als de reparateur twee keer zo snel zou repareren, blijkt al een heel stuk moeilijker te beantwoorden. Dat de reparatiesnelheid invloed heeft op de verblijftijd van een product moge duidelijk zijn: indien

een reparateur heel traag zou zijn, wordt in gemiddelde zin de tijd dat een machine kapot blijft veel groter, en daarmee ook de verblijftijd van de producten. Het daadwerkelijke verband tussen deze twee zaken blijkt echter zeer moeilijk aan te geven. De reparateur heeft immers geen directe rol ten aanzien van de producten, maar beïnvloedt de verblijftijd slechts indirect via de duale rol van de machine als klant van de reparateur en bediende van de producten. Als gevolg hiervan ontstaan er interacties tussen de twee lagen, die de wiskundige analyse van het netwerk bemoeilijken.

Daarnaast komt met de aanwezigheid van een reparateur (ofwel de tweede laag van het netwerk) een andere vervelende bijkomstigheid om de hoek kijken. In het geval beide machines niet kapot kunnen gaan, en geen reparateur nodig is, zijn de verblijftijden van producten in de bovenste wachtrij van de eerste laag compleet onafhankelijk van die van de producten in de onderste wachtrij. De machines bedienen de producten immers onafhankelijk van elkaar in de aan hen toegevoerde wachtrijen. In het geval dat de machines kapot kunnen gaan, ligt dit echter anders. Als beide machines tegelijkertijd kapot zijn, kan de reparateur maar aan de reparatie van één machine tegelijk werken, waarop de andere machine zal moeten wachten. Dit maakt dat de zogenaamde 'downtijden' van beide machines statistische correlatie gaan vertonen. Immers, als de reparatie van de ene machine lang duurt, zal de andere machine langer moeten wachten dan gebruikelijk. Dit effect in de tweede laag van het netwerk vertaalt zich via de duale rol van de machines door naar de eerste laag, waardoor verblijftijden van producten in verschillende wachtrijen statistisch gecorreleerd raken. Opnieuw is het precieze verband in algemene zin lastig te bepalen.

Limietregimes

Ondanks de aanwezigheid van gelaagde structuren in praktijksituaties, staat de aandacht die in de vakliteratuur besteed is aan gelaagde wachtrijnetwerken in schril contrast met die voor de niet-gelaagde varianten. Het lijkt geen twijfel dat dit mede een gevolg is van de hierboven beschreven uitdagingen die gelaagdheid met zich meebrengt. Daarentegen leggen wachtrijnetwerken die zich voordoen in hedendaagse applicaties steeds meer een gelaagd karakter aan de dag, waardoor een prestatie-analyse van gelaagde systemen noodzakelijk is.

Om voor het voorbeeld van het productiesysteem

toch een prestatie-analyse mogelijk te maken, zou men het model kunnen bestuderen in bepaalde limietregimes. Het blijkt in het geval dat de machines nauwelijks producten aangeboden krijgen (het *light-traffic* regime in vaktermen) het mogelijk is om uitdrukkingen voor belangrijke prestatie-indicatoren af te leiden zoals gemiddelde verblijftijden van producten of gemiddelde rijlengtes in de eerste laag. Hetzelfde blijkt het geval wanneer de binnenkomende te verwerken productiestromen juist zo groot zijn dat de machines op het punt staan overbelast te raken (het *heavy-traffic* regime). In de praktijk ligt de intensiteit van deze stromen ergens tussen deze regimes in, zodat een interpolatie tussen de limietresultaten van beide regimes zeer precieze benaderingen voor de werkelijke prestatie-indicatoren kan opleveren (Dorsman et al., 2013). Met deze en andere benaderingen is het mogelijk om grip te krijgen op de gevolgen van gelaagdheid in netwerken en op basis hiervan antwoorden te formuleren op belangrijke vragen zoals 'Welke machine kan de reparateur het beste repareren, als beide machines kapot zijn?'. Het antwoord op deze vraag hangt af van de hoeveelheid producten die in elke wachtrij aanwezig zijn, en correcte implementatie van dit antwoord blijkt enorme verbeteringen teweeg te kunnen brengen in de doorstroom van de producten in het productiesysteem (Dorsman et al., 2015).

Het onderzoeksterrein van gelaagde wachtrijnetwerken vormt een belangrijk en tot dusver vrijwel onontgonnen gebied. De hierboven genoemde eerste resultaten hebben de potentie een basis te vormen voor de prestatie-analyse en optimalisatie van tal van grotere en gelijksoortige netwerken van gelaagde wachtrijen.

LITERATUUR

- Dorsman, J.L., van der Mei, R.D., & Vlasiov, M. (2013). Analysis of a two-layered network by means of the power-series algorithm. *Performance Evaluation*, 70, 1072–1089.
- Dorsman, J.L., Bhulai, S., & Vlasiov, M. (2015). Dynamic server assignment in an extended machine-repair model. *IIE Transactions*, 47, 392–413.
- Jackson, J.R. (1957). Networks of waiting lines. *Operations Research*, 5, 518–521.

JAN-PIETER DORSMAN promoveerde op 17 februari 2015 aan de Technische Universiteit Eindhoven op het proefschrift 'Layered Queueing Networks: Performance Modelling, Analysis and Optimisation' onder begeleiding van Onno Boxma, Rob van der Mei en Maria Vlasiov. Voor dit proefschrift kreeg hij een eervolle vermelding in het kader van de VVS+OR W.R. van Zwet Thesis Award 2015. Momenteel is hij als postdoc verbonden aan de Universiteit van Amsterdam en Universiteit Leiden. E-mail: j.l.dorsman@uva.nl

Loterijen onder vuur

In een eerdere column in *STAtOR* (2016) beschreef ik hoe ongelooflijk klein de kans is om de jackpot in de lotto te winnen. Deelname aan loterijen komt in de meeste gevallen neer op vrijwillig belasting betalen. Maar niet elke deelnemer aan de loterij laat zich als een lam naar de slachtbank leiden. En daar wil ik in deze column op ingaan.

Als in de lotto de jackpot zodanig gegroeid is dat de verwachte uitbetaling op één rijtje groter dan de kosten van het rijtje, dan kan het lonend zijn om alle mogelijke rijtjes te kopen. Dit is wat sommige grote syndicaten met succes gedaan hebben. Het beroemdste geval is het syndicaat dat de wiskundige Stefan Mandel had opgezet. Met zijn reputatie als lottobreker had hij in Australië een consortium van 2500 investeerders gevormd die elk 4000 dollars hadden ingelegd. In 1992 koos hij als doelwit de Virginia lotto in de USA waarvan het deelname reglement het syndicaat ruimte bood om massaal loten in te kopen. In deze lotto worden 6 getallen uit de getallen 1 tot en met 44 getrokken en bij 6 goed win je de jackpot. Bij deze lotto zijn ongeveer 7 miljoen combinaties van 6 getallen mogelijk. In februari 1992 stond de jackpot op 27 miljoen dollars. Mandel had deze lotto al enige tijd gevolgd en de zaken minutieus voorbereid. Vooraf had Mandel in Australië om en nabij de 1,4 miljoen formulieren laten afdrukken met op elk formulier vijf ingevulde rijtjes zodat precies alle mogelijke combinaties gedekt waren (een efficiënt wiskundig algoritme had Mandel daartoe ontwikkeld). Deze formulieren werden verscheept naar een compagnon in Virginia die met militaire precisie een team had gevormd waarbij elk teamlid een aantal verkooppunten van de lotto en stapels in te leveren formulieren toegewezen kreeg. Het lukte het consortium niet alle 7 miljoen combinaties in te leveren, maar wel ongeveer 5 miljoen. Het consortium had het geluk dat er geen run was van andere deelnemers op de bewuste trekking zodat de kans vrij groot was dat bij winnen van de jackpot door het consortium de prijs niet met anderen gedeeld hoefde te worden. Hoe groot is deze kans? Wij kunnen alleen een idee van de grootte geven. Daartoe gaan we uit van 2 miljoen rijtjes ingevuld door de andere deelnemers. De kans dat de jackpot op een rijtje van het consortium valt is ongeveer $5/7$. De kans dat k andere personen met hun random gekozen getallen de jackpot winnen, kan benaderd worden door

de Poisson kans $e^{-\lambda} \lambda^k / k!$ met $\lambda = 2/7$. Deze Poisson kans heeft de waarde 0,7515 voor $k = 0$, de waarde 0,2147 voor $k = 1$ en de waarde 0,0307 voor $k = 2$. Dus de kans was ongeveer $5/7 \times 0,7515 \approx 0,537$ dat het consortium de jackpot als enige zou winnen. Dit gebeurde ook. Het consortium kreeg niet zonder slag of stoot het bedrag uitgekeerd, maar uiteindelijk maakte het consortium een stevige winst op de toch wel riskante investering. Het was ook het laatste kunstje van Stefan Mandel. Tegenwoordig woont deze van oorsprong Roemeense wiskundige op een tropisch eiland in de Stille Zuidzee.

In 2004 werd in de Amerikaanse staat Massachusetts de Cash WinFall loterij geïntroduceerd. Deze loterij werd ingevoerd omdat bij de vorige loterij de jackpot zelden gewonnen werd met als gevolg dat de deelname aan de loterij sterk terugliep. De Cash Winfall loterij waarin voor de jackpot nodig was alle zes goed van de zes getrokken getallen uit 1 tot en met 46, had het bijzondere kenmerk dat de gehele jackpot toegevoegd werd aan de prijzen voor 5 goed, 4 goed, en 3 goed als de jackpot niet viel en boven de 2 miljoen dollar stond. Deze prijzen van \$4000, \$150 en \$5 gingen dan met een factor vijf tot acht omhoog. Van dit gegeven en het feit dat de jackpot niet viel in bijna 99% van de tijd werd door enkele syndicaten slim gebruik gemaakt door een groot aantal loten te kopen op het moment dat een *roll-down* aanstaande was. Stel eens dat een syndicaat \$400.000 heeft geïnvesteerd in 200 duizend random ingevulde rijtjes van \$2 per stuk, terwijl de jackpot niet gevallen is en boven de twee miljoen dollar staat. Wat is dan de verwachte opbrengst voor het syndicaat als van te voren was aangekondigd dat een rijtje met 5 goed \$25.000 oplevert, een rijtje met 4 goed \$925 en een rijtje met 3 goed \$27,50? De verwachtingswaarden van het aantal rijtjes met 5 goed, 4 goed en 3 goed onder de voorwaarde van geen enkel rijtje met 6 goed zijn 5,12447, 249,8180 en 4219,149, zoals met het hypergeometrische kansmodel berekend kan worden. Bijgevolg is de verwachte opbrengst voor het syndicaat gelijk aan $5,12447 \times 25.000 + 249,8180 \times 925 + 4219,149 \times 27,50 = 475.220$ dollar. De kansverdeling van de totale opbrengst kan benaderd worden door de aantallen rijtjes met 5 goed, 4 goed, en 3 goed te modelleren als onafhankelijke Poisson variabelen, hetgeen leidt tot een benadering van \$58.479 voor de standaarddeviatie



Loterijtrekking in Kanton, circa 1910

van de opbrengst en vervolgens met de normale verdeling tot een benadering van 9,9% voor de kans dat het syndicaat de investering van \$400.000 niet terugverdient. Drie syndicaten waaronder een groep studenten van MIT wisten miljoenen dollars te winnen door slim gebruiken te maken van de *roll-down* eigenschap van de loterij, in feite profiterend van de opbouw van de jackpot door anderen. In 2011 kwamen de syndicaten negatief in de publiciteit en kort daarna werd de loterij opgeheven.

Loterijen en bedrog zijn onlosmakelijk met elkaar verbonden. Je hoeft maar even te 'koekelen' om voorbeelden te vinden. Zo claimde onlangs een Engelse oma de hoofdprijs van 33 miljoen Britse ponden in de Engelse National Lottery. Lange tijd had de winnaar van de hoofdprijs zich nog niet gemeld en toen kwam deze oma met het verhaal dat haar lot in de wasmachine had meegedraaid waardoor slechts enkele van de winnende getallen op haar lot nog zichtbaar waren. Deze claim kreeg veel publiciteit in de media, maar kort nadat dit verhaal naar buiten was gekomen meldde zich iemand met het echte winnende lot en werd oma door de loterij aangeklaagd voor bedrog.

Het meest beruchte geval van loterijfraude, bekend als de Triple Six Fix vond plaats in 1980 in Pennsylvania. In de betreffende loterij werden tijdens een TV show in drie machines ballen met de nummers 0 tot met 9 met een luchtblazer door elkaar gehusseld totdat uiteindelijk drie ballen een vacuüm tube hadden bereikt, waarbij de volgorde waarin deze drie ballen verschenen van belang was. De groep fraudeurs stond onder leiding van Nick Perry die de presentator van de tv-show was. Perry had

witte latex verf in elk van de ballen geïnjecteerd behalve in de ballen met de nummers 4 en 6. Dit betekende dat alleen ballen met één van deze twee nummers de vacuüm tube konden bereiken. Vervolgens had de groep heel veel loten gekocht met de acht combinaties 666, 664, 646, 644, 466, 464, 446 en 444.

In de bewuste trekking was de winnende combinatie 666, vandaar de naam Triple Six Fix. Het syndicaat van fraudeurs kon niet lang genieten van de ongeveer 1,8 miljoen dollar die ze gewonnen hadden. In een bar hadden leden van het syndicaat een groot aantal loten gekocht met combinaties van de getallen 4 en 6 en daarna had één van de leden een telefoongesprek in het Grieks gevoerd. Dit was een medewerker van de bar opgevalen en die meldde dit aan het tv-station nadat hij gehoord had over twijfels van lokale bookmakers over de uitkomst van de trekking waarvoor zo ongebruikelijk veel loten met alleen de cijfers 4 en 6 verkocht waren. Het telefoongesprek kon worden getraceerd naar de tvstudio en daar viel al snel de verdenking op Nick Perry die Grieks sprak. De fraudeurs werden gearresteerd en kregen stevige gevangenisstraffen waaronder zeven jaar voor Perry. In 2000 verscheen over dit gebeuren de film *Lucky Numbers* met John Travolta in de hoofdrol. Het is niet bekend of Perry na zijn vrijlating deze film ooit heeft bezocht.

LITERATUUR

Tijms, H. (2016). Een gokje in de loterij. *STAtOR*, 17, 3, 32–33

HENK TIJMS is emeritus hoogleraar operations research aan de Vrije Universiteit en auteur van diverse leerboeken over operations research en kansrekening. E-mail: tijms@quicknet.nl

Claude Shannon in de Bell Telephone Laboratories met een elektronische muis met 'super'-geheugen. Na slechts één trainingsronde vindt de muis moeiteloos zijn weg in het doolhof. 1952. Foto: Hulton Archive

VAN HERACLITUS TOT SHANNON

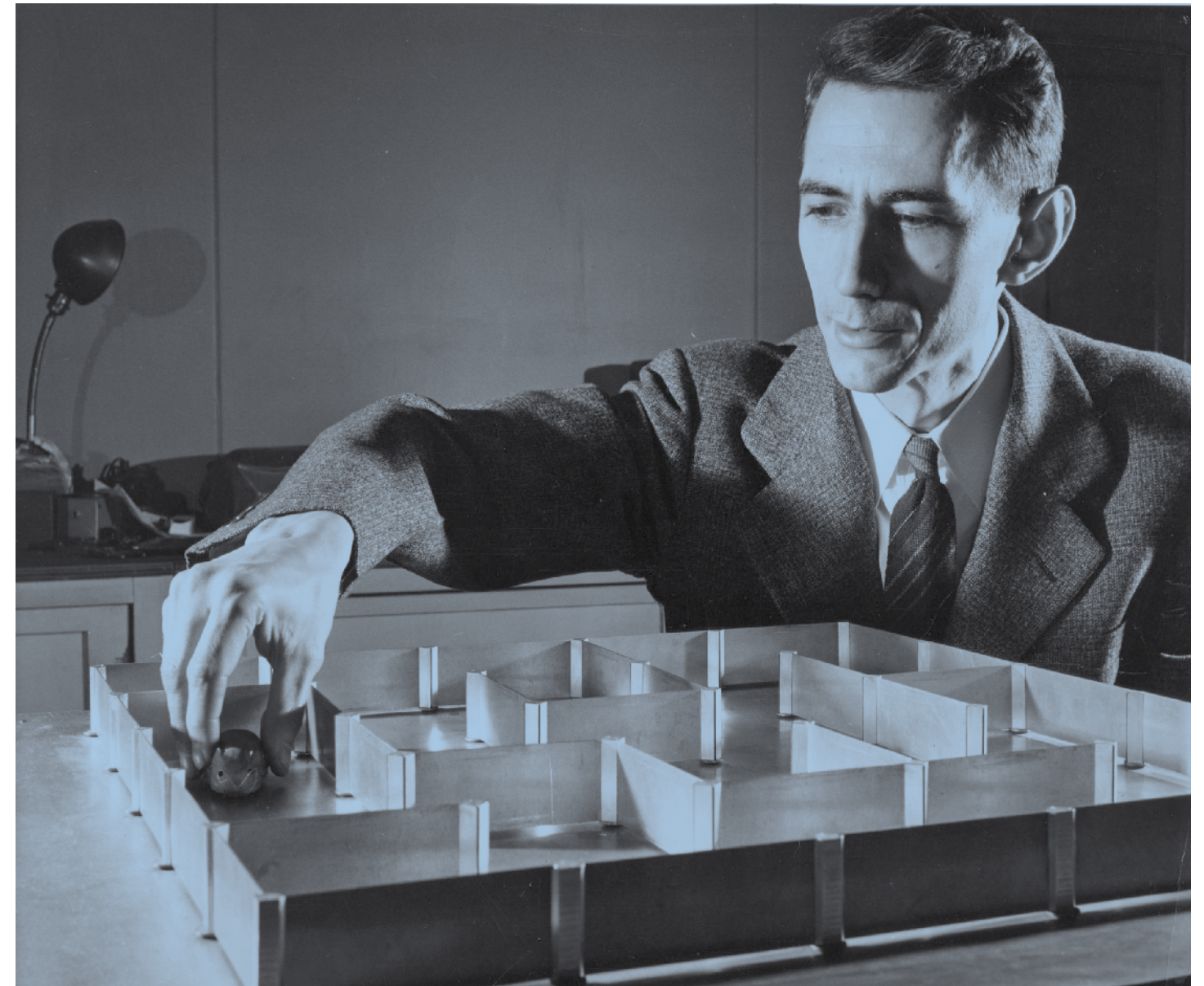
de fluwelen revolutie van data in context en flux

De Amerikaanse wiskundige, elektrotechnisch ingenieur en grondlegger van de wiskundige informatietheorie Claude Shannon wordt in overzichtswerken van de geschiedenis van kansrekening en statistiek dikwijls slechts summier besproken. Toch speelde hij, na de statisticus Karl Pearson en de fysicus Niels Bohr, een cruciale rol in de probabilisering van het wereldbeeld. Een historisch-filosofische blik op deze revolutionair tegen wil en dank.

RICHARD STARMANS

Een notoir geschilpunt onder wetenschapshistorici betreft de vraag of vooruitgang in de wetenschap zich evolutionair, dan wel revolutionair manifesteert. Gaat het om kleine, stapsgewijze aanpassingen en verbeteringen in een min of meer continu proces van kennisvermeerdering, waarbij de waarheid steeds dichter wordt benaderd? Of gaat een en ander gepaard met plotselinge, onvoorziene en radicale omwentelingen, breekpunten of discontinuïteiten? Thomas Kuhn (1922-1996), de na of naast Sir Karl Popper beroemdste wetenschapsfilosoof van de 20e eeuw, zette in zijn nog steeds invloedrijke *The Structure of Scientific Revolution* (1962) zijn paradigmatheorie uiteen, waarin het belang van 'revoluties' en discontinuïteit in de geschiedenis van de natuurwetenschappen wordt benadrukt. De Franse cultuurfilosoof

Michel Foucault (1926-1984) deed in *Les Mots et les Choses* (1966) met zijn 'archeologische' methode ruwweg hetzelfde voor de humaniora. Beiden lieten zien dat de geschiedenis cruciale gebeurtenissen of perioden van radicale omwentelingen kent, die de mens dwingen zijn eigen wezen, zijn positie in de kosmos, verantwoordelijkheden en zelfbeeld opnieuw te bezien. Beiden trachtten deze fundamentele cesuren aan te wijzen en te benoemen. Volgens de Italiaanse filosoof Luciano Floridi (1964) markeert de snelle opkomst van informatica en ICT zo'n cesuur. In *The Fourth Revolution, how the infosphere is reshaping human reality* (2014) postuleert hij een Vierde Revolutie, in gang gezet door de Britse wiskundige Alan Turing (1912-1954). Eerst maakte de Copernicaanse Revolutie ons duidelijk dat wij niet on-



beweeglijk in het centrum van het universum vertoeven. Daarna bewerkstelligde de Darwiniaanse Revolutie een drastische breuk met onze vertrouwde opvattingen over de menselijke soort en afkomst. Vervolgens leerde de Freudiaanse Revolutie ons dat wij allerminst volledig transparant voor onszelf zijn. De ICT, aldus Floridi, doet ons beseffen dat wij geen op zichzelf staande entiteiten zijn, maar informationele organismen (inforgs), die met andere (al dan niet artificiële) actoren een gemeenschappelijke leefwereld delen, welke uiteindelijk gemaakt is van informatie, de 'Infosphere'. Daarmee betoont Floridi zich een apologeet van een soort informationele metafysica, die vraagt om een herbezinning op vele traditionele metafysische, ethische en maatschappelijke denkbeelden. De huidige discussies over de risico's van Artificial

Intelligence (AI), Data Science en Big Data voeden de stelling dat een dergelijke herbezinning in volle gang is.

In dit korte essay plaatsen we slechts enkele kanttekeningen bij dit 'evolutie-revolutie debat'. De eerste betreft de conceptuele analyse van de tegenstelling zelf. Is deze wel goed gedefinieerd, coherent of slechts een saillant voorbeeld van een 'false dilemma'? Is het geen louter semantische kwestie, die afhangt van de striktheid waarmee verschillende auteurs beide begrippen definiëren? Is het bijvoorbeeld zinvol te spreken over de 'Wetenschappelijke Revolutie', terwijl het gaat om een periode die 100 tot 150 jaar duurde en die bovendien volgens historici als Pierre Duhem, Stevin Shapiro en James Hanman feitelijk een proces betrof dat al in de late Middeleeuwen een aanvang nam? En, als we ons

beperken tot kansrekening en statistiek: Is het zinvol om een ‘Probabilistische Revolutie’ te postuleren, zoals Lorenz Krüger cum suis doen in het lijvige, tweede-lige standaardwerk *The Probabilistic Revolution: Ideas in History* uit 1987? Vragen als deze laten zich bezwaarlijk ex cathedra beantwoorden, maar vergen op zijn minst nauwgezette historische case-studies.

De tweede kanttekening betreft een profylactische waarschuwing van de historicus van de statistiek Stephen Stigler (1941). In *Stigler’s Law of Eponymy* (1980) stelt hij onomwonden dat ‘no scientific theory is named after its original discoverer’ in weerwil van onze neiging om stellingen, deeltjes, constanten en zelfs theorieën, scholen of paradigma’s gewetensvol toe te schrijven aan één door ons bewonderde denker of wetenschapper. Dat de titel een ironische en destructieve zelfreferentie bevat laat onverlet dat Stigler, voortbouwend op het werk van de socioloog Robert K. Merton (1910-2003), één van de beperkingen of voetangels van de historiografie aanwijst. Hoe dan ook, als men al een revolutie meent te ontwaren, lijkt het een hachelijke onderneming deze aan één naam te verbinden, zoals Floridi – ongetwijfeld omwille van het mnemonische gemak – doet.

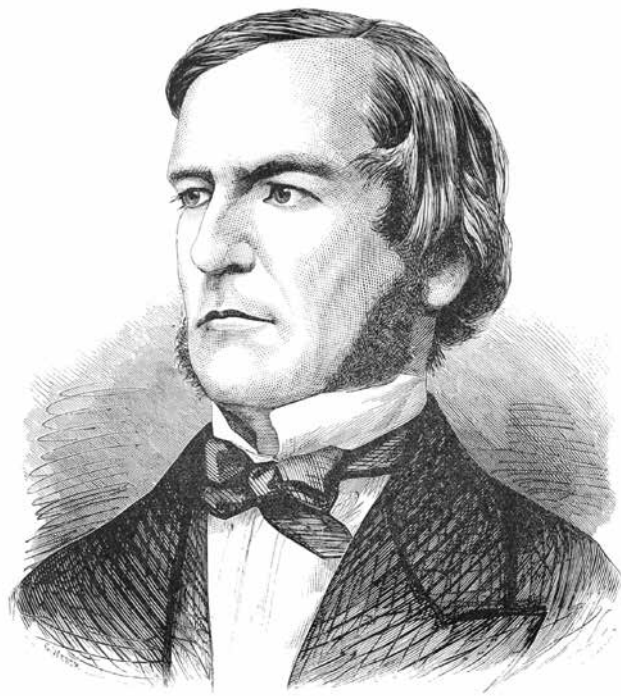
De derde kanttekening betreft een wedervraag, die hier als uitgangspunt wordt gehanteerd.

Als er al sprake is van een omwenteling, moet die dan ook beoogd en gewild zijn door de protagonisten, die zich opwerpen als hemelbestormers of zelfbenoemde titanen, of is zulk een revolutie wellicht ‘slechts’ een onbedoeld en mogelijk zelfs een deels onbegrepen en ongewenst effect, bijproduct of epifenomeen van inspanningen door gedreven zoekers naar orde en waarheid?

De ideeëngeschiedenis kent vanzelfsprekend talloze aanknopingspunten om deze vraag te beantwoorden. We beperken ons hier tot het werk van de Amerikaanse wiskundige, elektrotechnisch ingenieur en ‘probabilist par excellence’ Claude Shannon (1916–2001) wiens eerste centennial in 2016 mondiaal luister werd bijgezet.

Intermezzo: Max Planck

In het essay *De geest als matrix van de materie; Max Plancks revolutie tegen wil en dank* (Starmans, 2016)



George Boole. Bron: *Popular Science Monthly*, 1879, volume 17

wordt betoogd dat het werk van de Duitse natuurkundige Max Planck (1858–1947) een eenduidig antwoord toelaat op de voornoemde vraag. Planck was allerm minst het prototype van de hemelbestormer. Veeleer wilde hij een voortzetting van de bestaande traditie, maar bewerkstelligde een omwenteling die door hem noch bedoeld noch gewild was en die hij volgens sommigen ook niet geheel of pas betrekkelijk laat begreep. Plancks behoudzuchtige natuur bleek reeds op jonge leeftijd uit zijn reactie op het gedenkwaardige advies dat zijn latere leermeester Philipp von Jolly hem gaf. Deze raadde hem in 1874 aan vooral geen natuurkunde te gaan studeren, omdat in die discipline weinig nieuws meer te ontdekken viel. De jonge Planck koesterde geen ambitie deze achteraf notoir onjuiste inschatting van Von Jolly te weerspreken, maar riposteerde dat hij geenszins nieuwe dingen wilde ontdekken, doch de bestaande kennis beter wilde begrijpen om zo tot waar inzicht te komen. Ook later in zijn carrière manifesteerde zich deze attitude. Zo was hij fervent tegenstander van het atoommodel, dat een snelle opmars maakte en gaandeweg niet langer instrumentalistisch, als een wiskundige constructie werd geïnterpreteerd, maar realistisch; de atomen werden verondersteld daadwerkelijk te bestaan. Planck had zelf onderzoek gedaan naar de thermodynamica, maar was weinig gecharmeerd van de statistische aanpak van Ludwig Boltzmann, die de befaamde Tweede Hoofdwet van de Thermodynamica

als een waarschijnlijkheidswet opvatte. Daardoor werd de entropie het geheel van willekeurige bewegingen van deeltjes en was het niet langer zeker, maar slechts waarschijnlijk dat in een bepaald systeem de wanorde zou toenemen. Planck had eveneens een groot ontzag voor Maxwells klassieke elektrodynamica en wilde daaraan aanvankelijk niet tornen en voorkomen dat het denken over het wezen van licht zou worden teruggeworpen naar de periode van Huygens en Newton, toen ook al werd gediscussieerd over het golf- of deeltjeskarakter ervan. Om vergelijkbare redenen was hij aanvankelijk ook tegen Einsteins verklaring van het foto-elektrisch effect door middel van de introductie van fotonen, die nota bene op Plancks eigen werk waren gebaseerd!

Shannon en Boole

Om verschillende redenen heeft Claude Shannon een onwrikbare plaats in de wetenschapsgeschiedenis ingenomen. Het was de jonge Shannon, die als student wiskunde en elektrotechniek aan MIT op eenentwintigjarige leeftijd in zijn doctoraalscriptie *A Symbolic Analysis of Relay and Switching Circuits* (1937) zou wijzen op het belang van de Boolese algebra voor een systematisch ontwerp van elektronische schakelingen en circuits. Het ontwerpen en bouwen van telefooncentrales geschiedde indertijd vooral op basis van ad hoc principes, hetgeen door de sterke opkomst van de telefonie gaandeweg problematischer werd. Velen onderkenden de noodzaak van een *design-science*, maar Shannon zag als een van de eersten in dat Boole’s werk hiertoe een solide wiskundige grondslag bood. Zonder twijfel vormde Shannons thesis een preambule tot de komst van de elektronische computer.

Vanzelfsprekend kon Boole (1815–1864) dit alles niet voorzien toen hij in 1854 de beginselen van zijn formele methode publiceerde in de thans klassieke studie *An Investigation of the Laws of Thought, on Which are Founded the Mathematical Theories of Logic and Probabilities*. Dat hij desalniettemin wel degelijk een visionair was met verstrekkende ambities blijkt niet alleen uit de titel van voornoemd werk. Zoals uiteengezet in (Starmans, 2015) zou de sterk wijsgerig georiënteerde Boole de latere 20e eeuwse analytische filosofie, logica en wetenschapsfilo-

sophie ingrijpend beïnvloeden en laten zijn bijdragen aan de wetenschap zich als een triptiek duiden: formalisering van de (wijsgerige) logica in een algebraïsche calculus, integratie van logisch en probabilistisch redeneren, en dit alles gefundeerd in en corresponderend met een aanzet tot een computationele theorie van de menselijke geest.

Boole’s bijdragen aan het redeneren met onzekerheid, inductie en vooral kansrekening en statistiek waren uiteraard de jonge Claude Shannon niet ontgaan. Allereerst gaf Boole stevast een epistemische duiding aan het waarschijnlijkheidsbegrip; een kansuitspraak moet worden gezien tegen de achtergrond van de beschikbare kennis en de ‘verwachting’ van de persoon die zich aan de kansuitspraak committeert. Hij stond daarmee dicht bij Laplace’s epistemische kansbegrip, maar hield diens *principle of insufficient reason*, waarbij gebeurtenissen even waarschijnlijk worden geacht indien er onvoldoende reden is het tegengestelde aan te nemen. Kansuitspraken betroffen bij Boole logische relaties tussen proposities. Ofschoon vormen van redeneren met onzekerheid doorgaans inductief zijn en dus buiten de deductieve logica vallen, vormde dit voor Boole geen reden logica en waarschijnlijkheid principieel te scheiden of de laatste uit te sluiten van de ‘nieuwe’ symbolische orde. Dit alles geschiedde bijna 80 jaar voordat de Rus Andrej Kolmogorov (1903–1987) de kansrekening een solide axiomatische basis zou geven.

In Shannons doctoraalscriptie stond vooral de Boolese algebra centraal, maar in nagenoeg al zijn latere werk zou hij ook op probabilistisch gebied in de voetsporen van Boole treden en hem in veelzijdigheid naar de kroon steken. Anders dan Boole was Shannon geen filosoof pur sang of anderszins in de ‘grote orde der dingen’ geïnteresseerd. Evenmin was hij een zelfverklaarde revolutionair of unificator. Gedurende zijn gehele carrière hield Shannon zich primair bezig met praktische problemen op terreinen die hem fascineerden; variërend van jongleren, blackjack en speculeren op de beurs tot computerschaak en kunstmatige intelligentie, en variërend van genetica en speltheorie tot informatie- en communicatietheorie en cryptografie. Stevast onderzocht hij daarvan de wiskundige, in veel gevallen nadrukkelijk de probabilistische, aspecten, om vervolgens met een oplossing te komen die het belang van het oorspronkelijke

probleem verre oversteeg. In nagenoeg alle gevallen putte hij zich uit te benadrukken hoezeer hij schatplichtig was aan voorgangers als Boole en Boltzmann, tijdgenoten als John Tukey en John von Neumann en tevens nadrukkelijk aan zijn collega's van Bell Laboratories.

Informatie als communicatie

In weerwil van deze veelvuldigheid aan aandachtsgebieden staat Shannon in de wetenschapsgeschiedenis vooral geboekstaafd als de grondlegger van de informatietheorie, met name dankzij *A Mathematical Theory of Communication*, dat in 1948 in twee delen verscheen. Het inmiddels bijkans klassieke uitgangspunt ervan behelst de praktische vraag hoe een *signaal* (bericht) door een *zender* via een *kanaal* naar een *ontvanger* wordt verstuurd, hoe dit wordt *gecodeerd* en *gedecodeerd* en welke rol *ruis* en *redundantie* spelen. Diverse kwantitatieve, contextuele en dynamische aspecten van informatie werden op een elegante wiskundige, probabilistische wijze geanalyseerd c.q. van een maat voorzien: hoeveelheid, dichtheid, opslag, transport, waarde, variatie en spreiding, verwachting, onzekerheid, binaire codering et cetera. Alleen al deze bescheiden opsomming suggereert reeds dat Shannons informatie bovenal 'data in beweging' en 'data in context' behelst, waarover later meer. Het werk was daarnaast deels technisch van aard en redelijk toegankelijk voor niet-ingewijden.

Dat gold nog sterker voor de vereenvoudigde versie ervan die in 1963 onder de lichtelijk gewijzigde titel *The Mathematical Theory of Communication* het licht zag met Warren Weaver en Shannon als auteurs. Mede hierdoor rees de ster van Shannon in die wetenschappen waarin steeds vaker de fundamentele rol van informatie in het grondslagenonderzoek zichtbaar werd en waarin informatieprocessen vanuit verschillende invalshoeken werden bestudeerd of gesimuleerd. Voorbeelden zijn de hedendaagse wetenschapsfilosofie, informatica (Kolmogorov-complexiteit, MDL), AI (computational learning), fysica (Gibbs-entropie), biologie (DNA-codering) en economie (speltheorie). In vele disciplines had zich daarbij een probabilistische wending voorgedaan in de zin dat de voornaamste begrippen, methoden en zelfs het achterliggende wereldbeeld werden gebaseerd op of afhan-

kelijk waren van kansrekening en statistiek. Shannons werk kon daardoor uitgroeien tot een universele probabilistische theorie voor communicatie met toepassingen in de elektrotechniek, de biomedische wetenschappen, de natuurwetenschappen en ook de sociale wetenschappen en de humaniora. Ofschoon Shannon al tijdens zijn leven veelvuldig als unificator werd geroemd was hij zelf terughoudender: 'The word information has given different meanings by various writers in the general field of information theory. (...) It is hardly to be expected that a single concept of information would satisfactorily account for the numerous possible applications of this general field.' (Shannon, 1993). Dat hij daarin te bescheiden was moge uit het volgende blijken.

Informatie en entropie

Toen de digitale revolutie nog maar amper van start was gegaan, benadrukte menig zelfbewuste pionier van het ICT-tijdperk reeds het fundamentele karakter en zelfs de metafysische betekenis van informatie. Zo stipuleerde de Amerikaanse filosoof, wiskundige en grondlegger van de cybernetica Norbert Wiener (1894-1964) in een beroemd geworden citaat uit 1961 dat het begrip informatie fundamenteeler is dan materie en energie.

Information is information, not matter or energy. No materialism which does not admit this can survive the present day. Why does information matter? Because the structure of the universe – that is of matter and energy – is a function of the distribution and quantity of information. If matter and energy are randomly distributed through space, there is no structure to the universe; if there is useful information about the universe, then there are pockets of organized matter and energy within it, capable of doing work, such as being converted from one to the other. Information is the quality or property of bundles of matter or energy. (Wiener, 1961)

Dergelijke informationele beschrijvingen van de werkelijkheid zijn vanzelfsprekend koren op de molen van filosofen en wetenschappers die gebukt gaan onder het juk van een al te gemakkelijk materialisme of fysicaïstisch reductionisme. Ook vormden zij een opmaat tot de hedendaagse informationele metafysica, waarin de fysi-



Het graf van Ludwig Boltzmann in Wenen. Op de steen is de entropieformule $S = k \cdot \log W$ gegraveerd. Foto: Thomas Schneider

sche werkelijkheid wordt voorgesteld als een immense computer en fysische processen als berekeningen c.q. transitities/toestandsovergangen (Floridi, 2014).

De befaamde *It from bit*-doctrine van de fysicus John Archibald Wheeler (1911-2008) laat aan duidelijkheid evenmin te wensen over.

Otherwise put, every 'it' – every particle, every field of force, even the space-time continuum itself – derives its function, its meaning, its very existence entirely – even if in some contexts indirectly – from the apparatus-elicited answers to yes-or-no questions, binary choices, bits. 'It from bit' symbolizes the idea that every item of the physical world has at bottom – a very deep bottom, in most instances – an immaterial source and explanation; that which we call equipment-evoked responses; in short, that all things physical are information-theoretic in origin and that this is a participatory universe. (Wheeler, 1990)

Shannon maakte deze 'metafysica' mogelijk door de diepe verbondenheid tussen fysica en informatie in probabilistische termen te duiden en een nieuwe wijsgerige conceptie van informatie te bewerkstelligen, waarbij uiteraard zijn entropie-begrip een grote rol speelt. We moeten ons hier tot enkele facetten ervan beperken. Allereerst dient te worden opgemerkt dat zijn notie van informatie als data in context en data in flux op het eerste gezicht afbreuk lijkt te doen aan de dikwijls gehanteerde

triptiek data-informatie-kennis. Data verwijzen dan naar het fysische, akoestische signaal of het orthografische of anderszins symbolische teken ('een low level representatie'), informatie betreft het proces van betekenis toekennen, de semantiek, de interpretatie van deze signalen of tekens, terwijl kennis als *justified true belief* het hoogste niveau behelst, gericht op waarheid en geldigheid, bij uitstek het terrein van de epistemologie. Beperkt Shannon zich niet louter tot het eerste?

Feit is dat Shannon naar eigen zeggen aan het semantische begrip van informatie geen prioriteit wil geven. Hij stelt nadrukkelijk dat het fundamentele probleem van communicatie

... is that of reproducing at one point either exactly or approximately a message selected at another point. Frequently the messages have meaning; that is they refer to or are correlated according to some system with certain physical or conceptual entities. These semantic aspects of communication are irrelevant to the engineering problem. The significant aspect is that the actual message is one selected from a set of possible messages. The system must be designed to operate for each possible selection, not just the one which will actually be chosen since this is unknown at the time of design.

Centraal staan de 'statistical structure of the original message' en 'the nature of the final destination of the information'. (Shannon, 1948) Om informatie, inclusief dynamiek en context te begrijpen, zal het begrip (of aspecten ervan) als een fysische grootheid moeten worden beschouwd.

Can we define a quantity which will measure, in some sense, how much information is 'produced' by such a process, or better, at what rate information is produced?

In al haar eenvoud is deze formulering functioneel vaag; ze behelst diverse connotaties die verbonden zijn met data in flux en data in context. Eerst legt Shannon uit waarom een logaritmische functie noodzakelijk is bij kiezen van een eenheid om informatie te meten, daarna toont hij zich schatplichtig aan de statisticus John Tukey (1915-2000):

If the base 2 is used the resulting units may be called binary digits, or more briefly bits, a word suggested by J. W. Tukey. A device with two stable positions,

such as a relay or a flip-flop circuit, can store one bit of information.

De volgende stap is overbekend. Gegeven een verzameling mogelijke gebeurtenissen, waarvan slechts de kansen op optreden bekend zijn, kunnen we dan een maat H vinden of *how much 'choice' is involved in the selection of the event or of how uncertain we are of the outcome?*

Shannon toont aan dat de enige H die hieraan voldoet overeenkomt met

$$H = -K \sum_{i=1}^n p_i \log_2 p_i$$

Daarmee is Shannons entropie geïntroduceerd. Het blijkt een zeer intuïtieve notie met vele voor de hand liggende, gewenste eigenschappen: H is 0 wanneer een event zeker is en de overige kansen gelijk zijn aan 0; H is groter dan 0 wanneer alle kansen ongelijk 1 zijn; H neemt toe naarmate de kansen minder verschillen en wordt maximaal bij een uniforme verdeling, waarbij onzekerheid dus volledig is. De onzekerheid van een *joint event* is kleiner of gelijk aan de som van de kansen van de afzonderlijke gebeurtenissen. De conditionele entropie wordt een maat voor de gemiddelde onwetendheid over y wanneer de verschillende waarden van x bekend zijn, etc. De onzekerheid van y wordt nooit groter door kennis van x . Ze wordt verminderd, zolang x en y niet onafhankelijk zijn en anders is zij constant. Daarbij geldt onder meer dat:

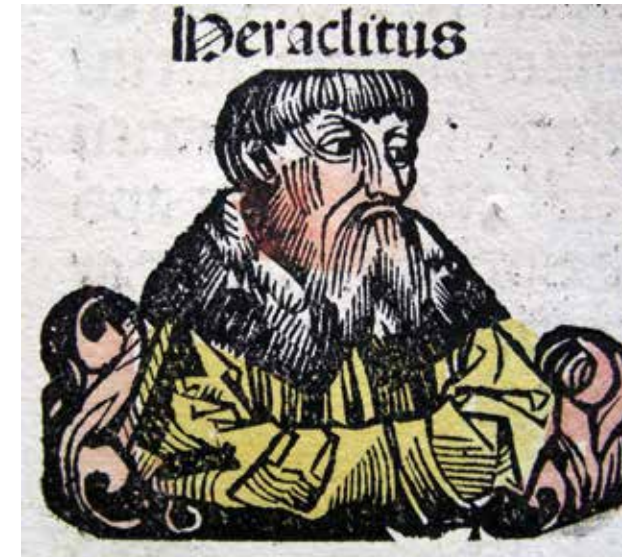
$$H(x) + H(y) \geq H(x, y) = H(x) + H_x(y)$$

Hierbij is $H_x(y)$ de voorwaardelijke entropie van y , gedefinieerd als de gemiddelde y voor elke waarde van x gewogen met de kans dat die waarde van x optreedt. Nadrukkelijk zoekt Shannon aansluiting bij en verwijst naar Ludwig Boltzmann (1844–1906). Met de opkomst van de kinetische gastheorie en statistische mechanica (Maxwell, Gibbs, Boltzmann), was de verbondenheid tussen fysica en statistiek irreversibel geworden. Shannon begreep met name de reikwijdte van het entropiebegrip in de thermodynamica. Ook hier krijgt een relatief eenvoudige formule pas reliëf in de context van het denken over warmte, energie en uiteraard de tweede hoofdwet van de thermodynamica. Deze kent diverse formuleringen/interpretaties en gevolgen, waardoor het

begrip vele connotaties heeft. De context behelst onder meer orde versus wanorde, gesloten versus open systemen, microscopische versus macroscopische toestanden, reversibele versus irreversibele processen, stabiele en instabiele evenwichten, omzetting van 'nuttige' of 'onnuttige' energie. Zo kan volgens de tweede hoofdwet warmte alleen stromen van een lichaam met hoge temperatuur naar een met lage temperatuur. Beschouwen we warmte als maat voor wanorde of entropie, dan moet de entropie van een geïsoleerd systeem dus toenemen. Ook een lokale of tijdelijke afname in een koud lichaam, waarin de atomen minder wanordelijk zijn dan de bewegende atomen in warme lichamen, wordt gecompenseerd door 'opwarming' en meer entropie in de omgeving, waardoor de netto-entropie onvermijdelijk toeneemt. Vooral buiten de fysica hebben overwegingen als deze soms aanleiding heeft gegeven tot zwaarmoedige bespiegelingen over het lot van de kosmos.

Shannon zag scherp in dat op analoge wijze zijn eenvoudige formule reliëf krijgt tegen de achtergrond van de vele facetten van data in context en flux. We beperken ons tot enkele aspecten, uiteraard onder de aanname dat de kansverdeling kan worden gespecificeerd. Zo beschouwd betekent entropie allereerst de gemiddelde hoeveelheid informatie per symbool die de zender verstrekt, maar is het tegelijkertijd een maat voor de onwetendheid van de ontvanger, preciezer gezegd, de gemiddelde hoeveelheid 'data deficit' die de zender produceert bij de ontvanger, zolang deze de uitkomst niet weet. Ook duidt entropie de informatiele potentie van de bron aan. De relatie met de thermodynamica is hier wellicht het meest dwingend. Entropie is een maat voor 'the amount of mixedupness in processes and systems bearing energy or information'. (Floridi, 2014). Hoe meer *randomness*, des te meer informatie kan door de bron of het device worden geproduceerd. Een perfect georganiseerd bericht heeft minder randomness en entropie, dus minder informatie. Hoe hoger de potentiële *randomness* van de data des te groter is het *vermogen* van de *source* om informatie te produceren. In de thermodynamica geldt dat hoe hoger de entropie, des te lager de beschikbaarheid/buikbaarheid van de energie. Zoals hoge entropie correspondeert met hoge *energy deficit*, zo correspondeert het in zekere zin ook met hoge *data deficit*.

De lijst met analogieën kan moeiteloos worden uit-



Heraclitus. Illustratie in het boek Liber Chronicarum. Nuremberg, 1493

gebreid en Shannons keuze voor het begrip entropie was dus geenszins een poging tot mystificatie of een ludieke pastiche, maar benadrukt de wezenlijke verbondenheid tussen natuurkunde en informatie, die essentieel probabilistisch is.

De probabilisering van het wereldbeeld

De verwevenheid van natuurkunde enerzijds en kansrekening en statistiek anderzijds is in historisch en wijsgeurig opzicht zeer dwingend en kan vanuit diverse invalshoeken worden bestudeerd of toegelicht. In (Starmans, 2014) wordt de verwantschap getoond vanuit het denken over onzekerheid, dat in de ideeëngeschiedenis een lange genealogie kent. De conceptie van onzekerheid is van oudsher verbonden met de klassieke vraag naar het wezen van variatie en verandering; verandering van plaats (beweging), ontstaan en vergaan, verandering van hoedanigheid en hoeveelheid. Alleen Heraclitus stelde dat deze zowel ontologisch als epistemisch niet alleen mogelijk, maar zelfs noodzakelijk waren, maar voor de meeste Grieken was zijn dynamische filosofie problematisch. In het kielzog van Parmenides effenden Plato, Aristoteles en zelfs de atomisten het pad voor het latere dominante substantiedenken. Het bestaan van verandering werd dikwijls ontkend, onmogelijk geacht of gecorrigeerd gereduceerd tot non-verandering. De ermee ver-

bonden conceptie van onzekerheid vormde gaandeweg een aporie, die zich niet liet opheffen of 'wegverklaren' en die uit arren moede dan maar moest worden betoeteld, beheerst en meetbaar gemaakt. Daarbij werd getracht onzekerheid te reduceren tot verwante begrippen als onbetrouwbaarheid, onvolmaaktheid, onvoorspelbaarheid, onvolledigheid en onbepaaldheid.

Onzekerheid doorliep een emancipatieproces, dat uitgerekend dankzij statistiek en natuurkunde en hun wisselwerking in een tweetal stappen werd voltooid. (Starmans, 2014) De eerste stap kan voor een aanzienlijk deel aan Karl Pearson (1857–1936) worden toegeschreven. Pioniers van de statistiek als Laplace en Quetelet trachtten al dan niet met behulp van 'foutentheorie' de weerbarstige verschijnselen in het keurslijf van de normale verdeling te dwingen. Pearson daarentegen deed recht aan variatie en verandering door deze niet te identificeren in *errors*, maar in de verschijnselen zelf (gecodeerd in data), en terug te voeren tot verschillende (klassen van) kansverdelingen. Hij zag in dat vele verschijnselen niet normaal, doch scheef waren verdeeld en met behulp van vier parameters (gemiddelde, variantie, scheefheid en welving) konden worden beschreven en geclassificeerd. De variabiliteit in de natuur manifesteerde zich in een puntenwolk van metingen en Pearson zocht naar het *best fitting* model, de functie die het best paste bij de data. Deze was veeleer een zuinige beschrijving hiervan dan een causaal onderliggend mechanisme waardoor deze data waren gegenereerd. Als eerste gaf hij daarmee kansverdelingen een volwaardige plaats in de wetenschap en zag de wereld op een niveau van abstractie waarbij data, variatie in data, datagenererende mechanismen en parameters van de kansverdelingen de werkelijkheid veeleer coderen en opbouwen en niet zozeer een (vermeende) fysische werkelijkheid representeren of afbeelden. Ruwweg was met Pearsons statistisch wiskundige benadering van variatie en verandering (en Fishers latere synthese van evolutie en mendeliaanse genetica) de eerste transformatie in het denken over onzekerheid een feit.

De tweede stap werd gezet door de Deense natuurkundige Niels Bohr (1885–1962), die het verbluffende sluitstuk van voornoemd emancipatieproces vormde en samen met Werner Heisenberg definitief onzekerheid als bouwsteen van de werkelijkheid introduceerde, vooral

door de Kopenhaagse interpretatie van de kwantummechanica. Een dimensie van onzekerheid die hierbij naar voren komt is die van onbepaaldheid, het gegeven dat een eigenschap of toestand (nog) niet vastligt of is gerealiseerd c.q. ook niet onafhankelijk van de waarneming is te bepalen. Bohr beschouwde waarnemer en experiment als een enkel samenhangend systeem; met het meten wordt het probleem vastgelegd. Een niet gemeten deeltje is als het ware onbepaald, heeft geen geschiedenis. De onzekerheid is wezenlijk, niet te reduceren en geen gevolg van een beperkt kenvermogen. Daarnaast, of wellicht daardoor, veroorzaakte Bohr onbedoeld grote filosofische verwarring. Zo wordt hij (ten onrechte) dikwijls beschouwd als een instrumentalist, die elke vorm van wetenschappelijk realisme onmogelijk had gemaakt. Ook zou hij het principe van intelligibiliteit hebben ondermijnd, onder meer door de natuurlijke categorie van oorzaak-gevolg relaties overboord te zetten en de poort te openen naar subjectivisme, holisme of oosters obscurantisme. De kloof tussen natuurwetenschap en filosofie werd hiermee verdiept, terwijl Bohr zich nu juist bekommerde om het begrip van het wetenschappelijke wereldbeeld; hij zocht een zinvolle interpretatie door in een conceptuele analyse onder meer de beperkingen van de op de macroscopische wereld gebaseerde natuurlijke taal te analyseren.

De transformaties van Pearson (variatie, verandering, kansverdeling) en Bohr (onbepaaldheid, 'ontologische' onzekerheid) vormen mijlpalen in het geschetste emancipatieproces van het denken over onzekerheid, de teloorgang van de aanschouwelijkheid van het wereldbeeld en het vestigen van een probabilistisch universum. Bij dit alles vormt evenwel de bijdrage van Claude Shannon toch een belangrijke *katharsis*, door de (herintroductie van) epistemische onzekerheid en doordat hij het primaat legde bij een notie van informatie, die geen epifenomeen is, reduceerbaar, elimineerbaar of obsoleet. Belangrijke aspecten van de menselijke conditie vormen nu eenmaal de gesitueerdheid, intentionaliteit (gerichtheid, *aboutness*) en betekenistoekenning. Signalen, tekens, symbolen vragen interpretatie: het begrijpen van de natuur, sacrale teksten, openbaringen, poëzie, de geschiedenis, verborgen codes, et cetera. Shannon had weliswaar niet het oogmerk de traditionele notie van semantische informatie te 'coveren', maar zijn data in

context en flux – waarvan de belangrijkste aspecten probabilistisch zijn en management van onzekerheid betreffen – vormen nu juist de basis voor een betekenisvolle theorie over informatie waar filosofen lange tijd naar zochten. Het proces van betekenistoekenning is alleen niet louter statisch zoals bij een modeltheoretische of denotationele semantiek, maar vooral dynamisch, zoals hierboven kort geschetst. En dit op een manier die niet alleen de metafysische citaten van Wiener, Wheeler en Floridi begrijpelijk maakt, maar ook is terug te voeren tot het oude denken over variatie en verandering van Heraclitus, zonder welke informatie niet eens mogelijk of denkbaar is. Shannon ondermijnde de triptiek data-informatie-kennis niet, maar geeft het denken over informatie een impuls die vele wijsgerige problemen oploste. Zijn aanpak is paradigmatisch gebleken voor veel filosofen, variërend van Dretske's *Knowledge and the Flow of Information* (1981) tot de Amerikaanse filosoof Andrea Scarantino die in zijn recente 'Information as a Probabilistic Difference Maker' (*Australian Journal of Philosophy*, 2015) aansluiting zoekt bij de pragmatische traditie en bayesiaanse confirmatietheorie, maar volledig is gefundeerd in Shannons benadering. De lijst kan moeiteloos worden uitgebreid.

Epiloog

Shannon was evenals Planck geen revolutionair, maar anders dan de fysicus minder geïnteresseerd in de grote orde der dingen. Vanuit praktische problemen vond hij stevast oplossingen die veel verder reikten dan het oorspronkelijke probleem. Hij werd een unificator tegen wil en dank, die voortbouwde op vele grote resultaten, waarvan hij de implicaties beter begreep dat de geestelijke vaders zelf. Omdat zijn probabilistische en dynamische kijk op informatie zulk een universele toepassing vond, en vernieuwing op allerlei vlakken mogelijk maakte moet tegen de geschetste achtergrond wellicht over een fluwelen revolutie worden gesproken. Gelet op de cruciale rol van Shannon in het probabilistische denken is het curieus dat in de geschiedenis van kansrekening en statistiek slechts een bescheiden plaats voor hem is ingeruimd en hij in overzichtswerken dikwijls slechts terloops en summier wordt vermeld. Als de geschiedenis

van data science, waarin de grenzen tussen inferentiële statistiek en informatica/intelligente data analyse minder stringent zijn, geschreven gaat worden zal hierop ongetwijfeld een correctie worden aangebracht.

LITERATUUR

- Dretske, F.I. (1999). *Knowledge and the flow of information*. Cambridge University Press.
- Floridi, L. (2014). *The Fourth Revolution, how the infosphere is reshaping human reality*. Oxford.
- Scarantino, A. (2015). Information as a probabilistic difference maker. *Australian Journal of Philosophy*, 93, 3
- Shannon, C.E. (1948). A mathematical theory of communication. *The Bell System Technical Journal*, 27, 379–423, 623–656.
- Shannon, C.E. (1993). The lattice theory of information. In: N.J.A. Sloane & A.D. Wyner (Eds.), *The collected papers of Claude E. Shannon*, eds. N.J.A. Sloane & A.D. Wyner. Los Alamos, CA: IEEE Computer Society Press.
- Starmans, R.J.C.M. (2015). Tussen Hobbes en Turing; George Boole en de Wetten van het Denken. In: *Filosofie*, 25, 6.
- Starmans, R.J.C.M. (2014). De weg naar Kopenhagen; Statistiek, natuurkunde en onzekerheid. *STATOR*, 15, 2.
- Starmans, R.J.C.M. (2016). De geest als matrix van de materie; Max Plancks revolutie tegen wil en dank. *Filosofie*, 26, 1.
- Wheeler, J.A. (1990). Information, physics, quantum: the search for links. In: W.H. Zureck, *Complexity, entropy and the physics of information*. Redwood City, CA: Addison Wesley.
- Wiener, N. (1961). *Cybernetics or control and communication in the animal and the machine*. Cambridge, MA: MIT Press.

RICHARD STARMANS is verbonden aan de Faculteit Bèta-wetenschappen (Department of Information and Computing Sciences) van de Universiteit Utrecht. Hij doet onderzoek op het snijvlak van filosofie, statistiek en informatica. E-mail: starmans@cs.uu.nl



In dit nummer geen bijdrage in de rubriek Klassiekers. We hebben enkele verzoeken uitstaan, maar graag nodigen we alle lezers uit om een bijdrage te leveren voor deze rubriek. Vertel aan uw medelezers welk boek of artikel voor u een beslissende rol heeft gespeeld in uw loopbaan. U kunt contact opnemen met Richard Starmans <Starmans@cs.uu.nl>; hij vertelt u graag meer over deze serie artikelen.

AGENDA

10–12 juli 2017, Vrije Universiteit Amsterdam
EURO Working Group Vehicle Routing and Logistics

Op de 6de conferentie van de EURO Working Group op het gebied van Vehicle Routing and Logistics (VeRoLog) worden de plenaire sessies verzorgd door prof. Guy Desaulniers (Polytechnique Montreal) en prof. David Pisinger (Technical University of Denmark). Verder zijn er de paperpresentaties en geven dr. Gunes Erdogan (University of Bath) en prof. Michael Schneider (RWTH Aachen University) praktische sessies over het oplossen van *vehicle routing*-problemen. Tijdens de conferentie eindigt de VeRoLog *solver challenge* (georganiseerd door ORTEC) waarin kwaliteitscontrole bij melkveehouders wordt geoptimaliseerd. Meer informatie over de conferentie en de *solver challenge* is te vinden op <https://verolog2017.sciencesconf.org/>.

April 12–13, 2017, Utrecht
11th International Multilevel Conference

The Conference on Multilevel Analysis will be about all aspects of statistical multilevel analysis: innovative applications, theory, software, and methodology. The conference will be in an informal style, with much room for discussion. Invited speakers are professors Ellen Hamaker and Laura Stapleton. Further information on abstract submission and registration can be found at <http://multilevel.fss.uu.nl/>

April 24–26, 2017, Hasselt University, Campus Diepenbeek
6th Channel Network Conference
of the International Biometric Society (IBS)

The IBS biennial conference is organized by the IBS regions Belgium, France, Great-Britain/Ireland and the Netherlands. This conference aims at gathering biostatisticians to discuss the newest statistical methodology for the analysis of biological and medical data. Usually, around 150 researchers from both academia and industry participate in the conference. We are proud that Peter Diggle and Rebecca DerSimonian have agreed to be the keynote speakers. In addition, we will have three invited sessions on: High-dimensional Bayesian variable selection by Alex Lewin; Analysis of human growth by Sophie Swinkels; and Neuroimaging by John Aston. Further information: www.uhasselt.be/channel-network-conference-2017

HET ONGEPUBLICEERDE BOEK VAN LAJOS TAKÁCS

ONNO BOXMA

Dit is het verhaal van een nooit gepubliceerd boek van de Hongaarse wiskundige Lajos Takács. Takács stierf op 4 december 2015 in Cleveland Heights, Ohio. Hij was een briljante wiskundige, die belangrijke bijdragen heeft geleverd aan de theorie van de stochastische processen, in onderwerpen als puntprocessen, ballot-stellingen, stochastische wandelingen, random grafen, vertakingsprocessen, fluctuatietheorie en wachtrijtheorie. Er staan 225 publicaties op zijn naam, waaronder zes boeken.

Dit verhaal is echter geen necrologie; daarvoor zij verwezen naar de publicaties van Boxma & Zacks (2016) en Haghighi (2015). Het is ook geen biografie of uitgebreide beschrijving van zijn werk; zie daarvoor het artikel van Bingham (1994) en dat van Dshalalow & Syski (1994) of het autobiografische verhaal (Takács, 1986). Het is ook geen lofzang op zijn boeken, al zou daar op zich alle reden toe zijn.

Zo schreef STAOR-columnist Henk Tijms mij: 'Zijn prachtige boek *Introduction to the Theory of Queues* (Takács, 1962) is het boek dat de meeste invloed op mij heeft gehad. Heel veel heb ik geleerd uit dit boekje over probabilistisch redeneren, daarbij werkend met conditionele kansen en conditionele verwachtingen. Takács was wat *applied probability* betreft zijn tijd ver vooruit.'

Waar dit verhaal wel over gaat is een ander boek van Takács, één dat nooit gepubliceerd is.

Op 23 mei 2015 schreef Takács een e-mail aan zijn oudleerling Haghighi, die de volgende opvallende passage bevatte (2015, p. 657–658).

'An additional note: in July 1973 I completed my book *Theory of Random Fluctuations*, which unfortunately still is in manuscript form. While I was working on the book, John Wiley was so interested that they offered me a contract before the book was finished. I did not like to sign a contract until my work was ready for the press. The finished manuscript turned out to be 1600 pages which Wiley considered too big for a book. They offered to publish it in lecture note form, which was unacceptable to me. I was also unwilling to shorten the MS, so it is still unpublished.'

In de contacten met Dalma, de weduwe van Takács, rond het schrijven van de necrologie (Boxma & Zacks, 2016) kwam naar voren dat zij het erg zou toejuichen als de stochastiek-gemeenschap alsnog kennis zou kunnen nemen van dit nooit gepubliceerde manuscript van, inderdaad, 1600 getypte pagina's. In samenspraak met Remco van der Hofstad, wetenschappelijk directeur van Eurandom, is toen besloten het manuscript op te nemen in de rapportenserie van Eurandom. Remco en ik bestudeerden eerst elk een hoofdstuk dat dicht bij ons onderzoeksgebied lag. We waren beiden onder de indruk van het materiaal. Ja, geda-

teerd, en getikt op een ouderwetse elektrische schrijfmachine; maar ook wiskundig diep, zorgvuldig opgeschreven, een schat aan materiaal, en met een enorm aantal interessante verwijzingen.

Ik bekeek hoofdstuk X, ge-



titeld *Queuing, Risk and Storage Processes*. 128 pagina's, 471 verwijzingen, met opvallend veel aandacht voor het Lévy-proces, een generalisatie van de Brownse beweging die tegenwoordig veel meer dan in 1973 in de belangstelling staat. Het gedeelte over wachtrijtheorie heeft een flinke overlap met de hoofdstukken I.6.6 en II.5 van Cohen's magnum opus *The Single Server Queue*, maar is desondanks heel waardevol.

Patty Koorn (Eurandom) heeft ervoor gezorgd dat het gehele boek, verdeeld over 12 pdf-files (voor de tien hoofdstukken, de appendix en de oplossingen van vraagstukken – 126 pagina's!) te vinden is op de Eurandom website, en dat de files doorzoekbaar zijn.

Als stochastische processen u interesseren, dan raad ik u aan eens elektronisch door het boek te bladeren: www.eurandom.tue.nl/reports/index.html. Het is dan bijna onmogelijk niet diepe bewondering te voelen voor dit werk van een buitengewone wetenschapper – werk dat gedoemd leek door niemand gelezen te worden.

LITERATUUR

- Bingham, N.H. (1994). The work of Lajos Takács on probability theory. *Journal of Applied Probability*, 31A, 29–39.
- Boxma, O.J., & Zacks, S.H. (2016). Lajos Takács. *Queuing Systems*, 82, 1–4.
- Dshalalow, J.H., & Syski, R., (1994). Lajos Takács and his work. *Journal of Applied Mathematics and Stochastic Analysis*, 7, 215–237.
- Haghighi, A.M. (2015). Lajos Takács' life and contributions to combinatorics. *Applied Mathematics: An International Journal*, 10, 636–658.
- Takács, L. (1962). *Introduction to the theory of queues*. New York: Oxford University Press.
- Takács, L. (1986). Chance or determinism. In: J. Gani (ed.), *The craft of probabilistic modelling: A collection of personal accounts*. New York: Springer-Verlag, pp. 139–149.

ONNO BOXMA is hoogleraar Stochastische Besliskunde aan de Faculteit Wiskunde & Informatica van de Technische Universiteit Eindhoven. E-mail: o.j.boxma@tue.nl

FRED STEUTEL

column

Kansrekening in de VS

In de Verenigde Staten kwam de kansrekening pas in de jaren dertig van de grond. Dat was deels het gevolg van de Jodenvervolging in Europa: wie niet weg was, was 'gezien'. Een lijstje 'Amerikaanse' kanscorfeëen met hun affilatie en geboorteland. Ik heb verschillende van hen persoonlijk gekend.

Samuel Karlin, Stanford (Polen)

Kai Lai Chung, Stanford (China)

Frank Spitzer, Cornell (Oostenrijk)

William Feller, Princeton (Kroatië)

Samuel Kotz, George Washington University (China)

Eugene Lukacs, Bowling Green (Hongarije)

Mark Kac, Cornell (Polen)

Michel Loève, Berkeley (Syrië)

Joop Kemperman, Rochester (Nederland)

Harry Kesten, Cornell (Nederland)

Fred Steutel, Baltimore (Nederland).

Ik eindig met de laatste regels van het lied *Elements* van Tom Lehrer:

*These are the only ones of which the news
has come to Harvard,
And there may be many others, but they
haven't been discovered.*

FRED STEUTEL is emeritus hoogleraar kansrekening aan de TU Eindhoven. E-mail: fsteutel@xs4all.nl



SEXY DATA SCIENCE

DANIEL OBERSKI*

Laatst deed ik iets wat vast wel meer statistici af en toe doen: ik maakte me hardop zorgen over de replicatie-crisis. P-hacking, HARKing, data dredging, researcher degrees of freedom¹, ... nu we er goed over na beginnen te denken is het eigenlijk nog een wonder dat er soms wél iets gerepliceerd wordt, jammerde ik. Een data scientist die mij had aangehoord vroeg toen: 'Waarom gebruiken die sociale wetenschappers niet gewoon een *holdout sample*?' Ja, waarom eigenlijk niet. Dat zou best wel eens kunnen werken! Simpel, maar een beetje als moderne kunst: je had er zelf aan kunnen denken maar dat heb je niet gedaan.

Een omgekeerd voorbeeld. Eén van de beroemdste data scientists ter wereld, Sandy Pentland van MIT, publiceerde in 2015 een artikel in *Science* waarin hij en zijn co-auteurs lieten zien dat 'vier spatio-temporele datapunten genoeg zijn om 90% van de individuen te herleiden' *ondanks anonimiseren vooraf* (de Montjoye e.a.,

2015). Niet alleen *Science*, maar ook *Nature* en zelfs de *New York Times* kopten geschokt: *With a few bits of data, researchers re-identify 'anonymous' data*.² In alle commotie was echter toch ook een groep onderzoekers minder ondersteboven van de resultaten. Een groep statistici uit Tarragona, gespecialiseerd in *statistical disclosure control*.³ Zij publiceerden een commentaar in *Science* waarin ze lieten zien dat anonimiseren wél tot adequate bescherming kan leiden. Tenminste, als je gebruik maakt van de over de laatste 40 jaar vergaarde statistische kennis over dit onderwerp (Sánchez e.a., 2016).⁴

Wat deze voorbeelden illustreren is het ontstaan van een synthese – goedschiks of kwaadschiks – tussen de 'twee culturen' die Leo Breiman vijftien jaar geleden al omschreef (Breiman, 2001). *In one corner: DATA SCIENCE*⁵ – kampioen pragmatisch aanpakken en ad-hoc scriptjes schrijven die iets nuttigs doen, kenner van algoritmes, presentator van business cases, koning van Facebook

tot Twitter, enz. enz. met helden als Tukey, Breiman, Andrew Ng, en, laten we zeggen, Hemelrijk.⁶ *In the other corner: STATISTIEK* – kampioen dingen tot op de bodem uitzoeken, koning van kans en consistentie, baron van bias tot steekproefvariantie, markies van modellering, randomizeerder van experimenten en kiezer van priors, met helden als Tukey, Breiman, Laplace, Bayes, Pearson, Fisher, Neyman, Rubin, en, in de vergelijking blijvend, Van Dantzig en Van Zwet.⁷

Twee werelden

Vanzelfsprekend is deze tegenstelling gechargeerd en is er al van oudsher enige overlap. Amazon gebruikt bijvoorbeeld *matrix factorization* om nieuwe producten aan te raden, en daar deed Gifi ook al aan (Van Der Heijden & Van Buuren, 2016). Tukey is een held in beide culturen; met name zijn artikel 'The future of data analysis' wordt beschouwd als onderdeel van de geschiedenis van beide velden. En biostatistici zijn al langer bezig met 'grote' datasets, vooral met meer kolommen dan rijen. De statisticus Tibshirani stelde om dat probleem op te lossen de lasso voor en die techniek vindt weer gretig aftrek in *machine learning* (Hastie e.a., 2015).⁸ Later hebben statistici als Van de Geer zo beide velden kunnen verbeteren.

Toch zijn statistiek en data science vaak nog twee werelden. Maar wel werelden die elkaar steeds vaker vinden. Dat is logisch, want beide velden groeien als kool op de vele nieuw ontgonnen datavelden en we zijn allemaal bezig met dezelfde problemen: data verzamelen, datakwaliteit beoordelen, data beschrijven, voorspellingen doen, oorzaken en gevolgen uit elkaar halen, bedrijven, overheden, en personen adviseren, resultaten communiceren aan een breed publiek, etc. Of je dat nu 'statistiek' of 'data science' wilt noemen, het ligt voor de hand dat je iets kan leren van een ander die er (ook) goed in is. Zeker als die ander nét een andere blik heeft,

zoals bovenstaande voorbeelden wilden illustreren.

Statistiek en data science kunnen dan ook veel voor elkaar betekenen. Enerzijds kennen statistici veel algemene principes die nuttig kunnen zijn in het bredere gebied van de data science. Neem bijvoorbeeld het gebruik van Twitter bij het bestuderen van rampen. Klinkt mooi, maar zijn tweets tijdens een overstroming missing at random? Of neem de zorgen van een bedrijf als Booking.com over welke 'A/B-tests' ze moeten doen en hoe ze er honderden kunnen vergelijken. Kunnen ze misschien veel geld besparen met *fractional factorial designs*? Wat te doen als je een beslisboom wilt schatten op herhaalde metingen? Weet iemand daar een oplossing voor? En hoe kun je je klanten eigenlijk het beste vragen naar hun mening? Is vragenlijstontwerp arbitrair of heeft iemand daar als eens onderzoek naar gedaan?

Anderzijds heeft data science de statistiek ook veel te bieden. Er zijn veel categorieën maar ik noem er drie. Ten eerste een aanpak: die van data analyseerders als Tukey en Jan de Leeuw; het pragmatisme van goed genoeg en op tijd boven 'correct'. Ten tweede expertise in het aanboren van nieuwe databronnen, het omgaan met APIs, database management systemen, andere programmeertalen dan R en het creëren van kwalitatief hoogwaardige software met behulp van *unit testing*, *version control*, *staging*, goede documentatie en andere goede gewoonten. Ten derde maken technieken uit de wereld van *machine learning* en *data mining*: *regularization*, neurale netwerken, *support vector machines*, en het concept van *embeddings*, om maar wat voorbeelden te noemen, nog geen deel uit van het standaardarsenaal van veel statistici, maar kunnen vaak een nuttig gereedschap bieden voor hun taken.

Broodnodige kruisbestuiving

In Nederland zijn er al veel voorbeelden van samenwerking en, laten we eerlijk zijn, *rebranding*. Wie vroeger

statistisch consult gaf noemt zich nu vaak ook ‘data scientist’. Studenten van onze universiteiten vinden werk als data scientist bij overheden, bedrijven, en onderzoeksinstituten. Master-opleidingen zijn inhoudelijk aangepast en bevatten steeds meer data science. Daarbij zijn statistici, waaronder leden van de VvS+OR, nauw betrokken bij de data science-initiatieven van Nederlandse universiteiten, zoals die van Amsterdam, Leiden, Eindhoven, Tilburg, en (sinds kort) Utrecht. Dankzij dit alles is, in tegenstelling tot de situatie op andere plaatsen, de statistiek in Nederland vrij goed vertegenwoordigd in data science.

Toch kan er mijns inziens nog veel meer gebeuren, en dan vooral op het gebied van onderzoek. De eerste stap, *rebranding* en samenwerking, is gezet. De tweede, veel moeilijker, stap is het daadwerkelijk beter maken van beide vakgebieden door kruisbestuiving. En op dat gebied wil de VvS+OR een belangrijke bijdrage leveren.

Om die bijdrage te coördineren en data science en statistiek inhoudelijk bij elkaar te brengen hebben we daarom de spiksplinternieuwe Sectie Data Science opgericht. Onze doelstellingen zijn:

- Samenbrengen van mensen die met data science bezig zijn uit *alle* onderzoeksvelden, inclusief statistiek, OR, informatica, en toegepaste wetenschappen, met als doel het *wederzijds* uitwisselen van kennis en verbeteren van methoden;
- Verbreden VvS+OR door nieuwe aanwas uit andere vakgebieden;
- Profileren statistiek in Nederland via data science;
- Coördineren en organiseren van concrete activiteiten, zoals:
 - Hackathons (een eerste voorproefje zal plaatsvinden op de Dag voor Statistiek en Besliskunde op 23 maart);
 - Lezingen en webinars informatici, statistici, data scientists voor gemengd publiek (aftrap wordt gegeven door de lezing van Max Welling op 23 maart);
 - Andere activiteiten die de lezer in wil brengen!

De Sectie Data Science is actief betrokken bij de organisatie van de Dag voor Statistiek en Besliskunde. Het bestuur van de sectie bestaat uit Daniel Oberski (voorzitter), Katrijn Van Deun, Alessandro Di Bucchianico, Arend Oosterhoorn, en Maarten Joosen met ondersteuning van Erik-Jan van Kesteren. Zie ook de website <http://sectiedatascience.nl/>.

* Met dank aan Katrijn van Deun en Alessandro Di Bucchianico

NOTEN

1. <https://fivethirtyeight.com/features/science-isnt-broken>
2. <https://bits.blogs.nytimes.com/2015/01/29/with-a-few-bits-of-data-researchers-identify-anonymous-people/>
3. Voor lezers die niet bekend zijn met dat veld: het zijn de technieken die bijvoorbeeld het CBS al sinds jaar en dag gebruikt om vast te stellen of statistieken die zij maken niet herleid kunnen worden tot personen. Zie bijvoorbeeld Hundepool e.a. (2012).
4. Overigens kreeg de groep van Pentland het laatste woord (de Montjoye & Pentland, 2016) dat ze aanwendden om zich te beklagen over het gebrekkig begrip van de statistici van de *big*-heid van hun data. De discussie is dus verre van afgelopen!
5. Op de vraag of data science nu eigenlijk statistiek is of andersom of misschien wel geen van beide ga ik hier niet in. Ook zal ik sommige termen niet uit het Engels vertalen als ik denk dat dat verwarring op zou leveren.
6. Als Hemelrijk meer data had gehad. Dank aan Peter Grünwald voor deze suggestie.
7. Mijn oprechte excuses als uw voorkeur er niet tussen staat, of uw afkeur wel. Hopelijk geeft het nochtans een idee van de tegenstelling die door Breiman werd geschetst. (Ik laat me overigens graag corrigeren).
8. Kleine kanttekening hierbij: Breiman stelde als eerste een methode voor variabelenselectie voor, quasi gelijktijdig (één jaar eerder) met de lasso en dit is de nonnegative garrote.

LITERATUUR

- Breiman, L. (2001). Statistical Modeling: The Two Cultures (with comments and a rejoinder by the author). *Statistical Science: A Review Journal of the Institute of Mathematical Statistics*, 16(3), 199–231.
- De Montjoye, Y.-A., & Pentland, A.S. (2016). Response to Comment on ‘Unique in the shopping mall: On the re-identifiability of credit card metadata’. *Science*, 351(6279), 1274.
- De Montjoye, Y.-A., Radaelli, L., Singh, V.K., & Pentland, A.S. (2015). Unique in the shopping mall: On the re-identifiability of credit card metadata. *Science*, 347(6221), 536–539.
- Hastie, T., Tibshirani, R., & Wainwright, M. (2015). *Statistical Learning with Sparsity*. CRC Press.
- Hundepool, A., Domingo-Ferrer, J., Franconi, L., Giessing, S., Nordholt, E.S., Spicer, K., & De Wolf, P.-P. (2012). *Statistical Disclosure Control*. John Wiley & Sons.
- Sánchez, D., Martínez, S., & Domingo-Ferrer, J. (2016). Comment on ‘Unique in the shopping mall: On the re-identifiability of credit card metadata’. *Science*, 351(6279), 1274–1274.
- Van der Heijden, P.G.M., & Van Buuren, S. (2016). Looking back at the Gifi system of nonlinear multivariate analysis. *Journal of Statistical Software*, 73(4).

Daniel Oberski is Universitair Hoofddocent Data Science Methodologie bij de vakgroep Methodologie en Statistiek van de Universiteit Utrecht.
E-mail: d.l.oberski@uu.nl. Twitter: @DanielOberski

GERRIT STEMERDINK

column



Moeten mannen eerst worden genoemd en daarna vrouwen? Of is het andersom? Denk maar niet dat zoiets niet belangrijk is.

Ik neem geen standpunt in, al was het maar dat ik niet verwikkeld wil raken in welke feministische discussie dan ook. Ik neem slechts waar. Nog steeds is het in de westerse wereld gebruikelijk om mannen voorop te stellen. Zo heet mijn echtgenote J.C. Stermerdink-van de Wetering. Maar sinds enige tijd is de wet veranderd, ik vind dat volkomen terecht. Bij een huwelijk kan men nu kiezen hoe men wil heten. Als ik nu zou trouwen zou ik als man de naam Van de Wetering-Stermerdink mogen voeren.

In de wetenschap kan over de man/vrouw volgorde óók verschillend worden gedacht. Zelfs binnen een en hetzelfde vakgebied als de biologie! Daar heb ik een aardige ervaring mee gehad.

Een groot deel van de jaren tachtig van de vorige eeuw was ik werkzaam aan de Landbouwniversiteit Wageningen. Het was een heerlijke tijd, als statisticus bij het Rekencentrum kreeg ik alle vrijheid me in nationaal en internationaal verband bezig te houden met vergelijking van software op aspecten als rekennauwkeurigheid en gebruikersgemak. Ook werd ik veel geraadpleegd door onderzoekers uit alle mogelijke vakgebieden, de LUW was een heel brede universiteit, ondanks de beperking die de naam suggereert.

Ooit promoveerde daar een statisticus op een onderzoek naar variantiecomponenten in kruisingsschema's. Zoiets is belangrijk bij pogingen de eigenschappen van een ras te verbeteren, of dat nu melkkoeien of doerwtjes zijn. Gebruikelijk is dan dat er een stamboom wordt opgesteld en dat gekeken wordt in hoeverre de eigen-

schappen van de ouders in hun nakomelingen terug te vinden zijn. Er komen dan termen als *Fragaria virginiana* Mill. x *Fragaria chiloensis* L. aan te pas. Ik had het proefschrift met veel belangstelling gelezen, daarbij was me niets bijzonders opgevallen. Tenslotte had ik nooit enige aandacht besteed aan het verschijnsel man/vrouw-volgorde, ik vond dat eigenlijk niet relevant.

Maar Wageningen telde in die tijd een nogal activistische groep vrouwelijke wetenschappers. Dat ze zo activistisch waren zag ik als een logische reactie op een wereld die voor 80% door vaak bijzonder masculien denkende mannen werd bevolkt. Tijdens de promotie kwam een van deze dames met de opmerking dat het vreemd was dat in kruisingsschema's bij planten éérst de vrouw genoemd werd en dan de man, maar dat het bij dieren andersom was. Zij wilde weten wat de promovendus daarvan vond, ook al omdat het verwarrend was beide varianten in één proefschrift aan te treffen. Ik heb nooit kunnen achterhalen of deze vraag voorgekookt was, zoals veel vragen bij promoties. Maar het antwoord was grandioos. De letterlijke tekst weet ik niet meer, maar het kwam neer op: 'Ik vind het niet zo'n probleem, het is maar net wat men gewend is. We hebben hier te maken met twee separaat ontwikkelde tradities. Maar laat ik u gerust stellen, in tegenstelling tot de notatievolgorde staat tijdens de lijfelijke uitvoering van de kruising in de dierenwereld altijd het vrouwtje voor en het mannetje achter'. De aula barstte los in een geweldige lachbui, de rector magnificus had grote moeite het decorum van de plechtigheid te herstellen...

GERRIT STEMERDINK is eindredacteur van STAtOR.
E-mail: gjstermerdink@hotmail.com

IN MEMORIAM



Maarten van der Vlerk (1961–2016)

Op 9 oktober 2016 is prof. dr. Maarten H. van der Vlerk (54) overleden. Zijn overlijden was een schok voor de OR-gemeenschap en heeft zijn collega's en vrienden diep bedroefd. We zullen hem altijd herinneren om zijn warme persoonlijkheid, zijn eerlijkheid en zijn toewijding.

Maarten promoveerde in 1995 aan de Rijksuniversiteit Groningen op zijn proefschrift *Stochastic Programming with Integer Recourse*, een onderwerp waarin hij later een internationaal vooraanstaande rol zou gaan spelen. Na een periode aan CORE (Louvain-la-Neuve) keerde hij terug naar Groningen. In 1999 ontving hij een prestigieuze onderzoeksbeurs van de Koninklijke Nederlandse Akademie van Wetenschappen (KNAW).

Later werd hij benoemd tohoogleraar Stochastic Optimisation aan de Faculteit Economie en Bedrijfskunde (FEB) en tot opleidingsdirecteur van de bachelor- en masteropleidingen Econometrics & Operations Research en Econometrics, Operations Research & Actuarial Studies. Ook was hij lid van het bestuur van de LNMB, het Lande-

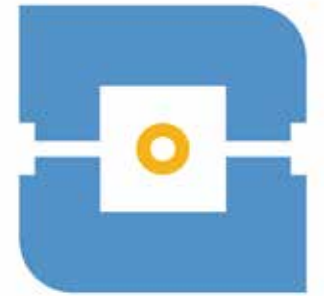
lijk Netwerk Mathematische Besliskunde, en was hij van 2004 tot 2007 voorzitter van COSP, het bestuur van de internationale gemeenschap op het gebied van Stochastisch Programmeren.

Maarten was niet alleen een briljante onderzoeker, internationaal erkend voor zijn werk in stochastisch programmeren, maar ook een geweldige docent met een gave voor het overbrengen van kennis. Hij was docent van het jaar aan de FEB in 2014. Zijn cursussen, Stochastic Programming en de Specialization Course Applied Operations Research, hoorden bijna elk jaar bij de top vijf van de faculteit. Vanaf de voorlichtingsdagen voor toekomstige studenten tot aan hun buluitreikingen, zijn toewijding aan studenten was voorbeeldig.

Onze gedachten gaan uit naar zijn collega's, vrienden en familie.

WIM KLEIN HANEVELD & WARD ROMEIJNDERS

THE YOUNG STATISTICIANS HAVE BEEN BUSY!



25 October: Statistics Café

YS wanted to start the year with a bang! So we started the year with a Statistics Café in Bar Beton, Utrecht. This was an evening for young statisticians, by young statisticians. Three Young Statisticians presented their work, for which we thank them:

- Sanne Blauw (de *Correspondent*) told us how data is changing investigative journalism, and how quantitative methods can be either misinterpreted or used to strengthen traditional news articles.
- Pieter Jongsma (Next Empire/Utrecht University) spoke about using neural networks, such as in the Prisma app, in ways in which they are not intended and using brute force to turn a picture into a painting.
- Michèle Nuijten (Meta Research group, Tilburg University) told us about the Statcheck project, which checks psychological papers for statistical errors. This program has sparked a huge discussion, such as in *Nature News*, and recently has been released as a web-app.

The evening had a very nice atmosphere, a lot of new contacts were made and the talks engaged us in interesting discussions. With over 50 persons attending it was a success.

16 November: Company visit at Delta Lloyd Group

Delta Lloyd Group is one of the largest insurance companies in the Netherlands, with brands like Delta Lloyd, ABN AMRO Insurance, BeFrank and OHRA. There were talks about risk management and events which are difficult to predict, such as natural catastrophes and market crashes, and about modelling Dutch mortality trends. After the talks, some employees and trainees showed up and there were some nice drinks and snacks. We found out that Delta Lloyd has some pretty nice and challenging traineeships.

21 November: Workshop Data Visualisation by Jan Willem Tulp

Jan Willem Tulp is an internationally renowned data experience designer, who recently visualized all trees on the planet for *Nature* magazine. Jan Willem gave us an interactive workshop of 4 hours, in which he discussed the design decisions that were involved in creating vi-

sual representations of data. Participants had to sketch their own data visualizations on paper and engage in active discussions. The workshop was well-visited.

8 December: Statistical Pubquiz: Statistics gone wrong

Huge scandals, badly used statistics, wrong news headers... Statistics has been having a hard time lately. We challenged the Young Statisticians to test their knowledge on statistical bloopers, but to make sure they were not too distressed we also amused them with some brain crackers, music, facial recognition problems and some other topics. The winners received a mug and a writing pillow, so that they can write down all their brilliant ideas when they wake up in the middle of the night. Three teams were fighting for the prize: Informative Priors, Survival of the Fitters, and Balkan Five. Survival of the Fitters won.

11 January: VvS+OR & Young Statisticians New Year's Drinks

To celebrate the new year, and to make sure that the junior and senior statisticians would mingle a bit the coming year, the board of YS and VvS+OR collaborated for the new year's drinks. The evening started out with drinks, and during the night the Stat-Olympics were started. The first 10 people to hand in a full *strippenkaart*, received a YS mug. 'Guess the number of candies in a pot', 'talk to someone you don't know', 'solve as many as possible problems in 2 minutes', were some of the assignments. It was fun!

UPCOMING EVENTS IN 2017

→ March > SAS Workshop

→ March > Lunch at the VvS+OR day

> (and party the day before)

→ April > Company Visit Google

→ May > Science Café

www.youngstatisticians.nl
contact@youngstatisticians.nl

