

STAtOR

Going Viral – Modeling Spread

Program Annual Meeting 2021
and the VVSOR Conference 2021

Bytes voor bites; Naar een wereld zonder honger

Vaccinatielogistiek essentieel om vaccinverspilling
en vertragingen te beperken

Met een beetje mathematische besliskunde
kom je een heel eind

Bezuinigen op verkiezingen met slimme wiskunde

Accumulation Bias; How to handle it ALL-IN

Hoe loop je een marathon?

Het balletje kan raar rollen

Met analytics de wereld verbeteren



STAtOR

Jaargang 22, nummer 1, maart 2021

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

Redactie

Joaquim Gromicho (hoofdredacteur), Annelieke Baller, Joep Burger, Kristiaan Glorie, Caroline Jagtenberg, Guus Luijben (eindredacteur), Kerry Malone, Richard Starmans, Gerrit Stermerdink (eindredacteur), Vanessa Torres van Grinsven en Sanne Willems. Vaste medewerkers: John Poppelaars, Gerard Sierksma en Henk Tijms.

Kopij en reacties richten aan

Prof. dr. J.A.S. Gromicho (hoofdredacteur), Universiteit van Amsterdam Faculteit Economie en Bedrijfskunde, Sectie Operations Management | Amsterdam Business School, Plantage Muidergracht 12, 1018 TV Amsterdam, j.a.s.gromicho@uva.nl

Bestuur van de VVSOR

Voorzitter: prof. dr. Fred van Eeuwijk, db@vvsor.nl; Secretaris: vacant, secretaris@vvsor.nl; Penningmeester: Judith ter Schure MSc, penningmeester@vvsor.nl; Bestuurscoördinator: Pieter Jongsma MSc, db@vvsor.nl; Algemeen bestuurslid: Thomas Wise MSc, db@vvsor.nl; Webmaster Eugenio Traini MSc: webmaster@vvsor.nl.

Voorzitters van de secties: prof. dr. ir. Mark van de Wiel (Biometrical Section); prof. dr. Albert Wagelmans (Section for Operations Research); dr. Eduard Belitser (Section Mathematical Statistics); prof. dr. Casper Albers (Social Sciences Section); dr. Michel van de Velden (Economics Section); dr. Daniel Oberski (Section Data Science); Jonas Haslbeck MSc (Young Statisticians); dr. Sanne Willems (Section Statistics Communication).

Leden- en abonnementenadministratie van de VVSOR

VVSOR, Maarsbergseweg 20, 3956 KW Leersum, admin@vvsor.nl. Raadpleeg onze website www.vvsor.nl over hoe u lid kunt worden van de VVSOR of een abonnement kunt nemen op STAtOR.

Voor advertenties

M. van Hootegem, hootegem@xs4all.nl
STAtOR verschijnt in maart, juni, september en december.

Ontwerp en opmaak

Pharos, Nijmegen

Uitgever

© Vereniging voor Statistiek en Operations Research
ISSN 1567-3383

Ons vak heeft impact

Natuurlijk weten we dat ons werk ertoe doet. Veel van ons onderzoek is direct op de praktijk gericht, maar ook puur theoretisch werk zien we al snel praktisch toegepast worden. Dit is de laatste maanden extra duidelijk geworden en bijna alle artikelen in dit nummer illustreren dat.

Het openingsartikel 'Bytes voor bites' van Marleen Balvert cs zouden we in brede kring moeten verspreiden: kijk eens wat er kan met statistiek en OR! In het Zero Hunger Lab van de Universiteit van Tilburg werken wetenschappers samen aan het wereldvoedselprobleem. Een betere reclame voor ons vak is nauwelijks denkbaar.

Wij nemen een artikel over van Jan Fransoo en Niels Agatz in *Medisch Contact* met de aanbeveling OR-specialisten te betrekken bij de distributie van COVID-19-vaccins. Tijdens het persklaar maken van STAtOR kwam het bericht dat Fransoo inmiddels met enkele collega's het RIVM hierin adviseert!

Heel praktisch is de toepassing die onze hoofdredacteur voor de school van zijn dochters maakte. Toen na de eerste lockdown leerlingen weer halve weken naar school mochten kon men moeiteloos de roosters zodanig inrichten dat kinderen uit hetzelfde gezin op dezelfde dagen werden ingedeeld.

Judith ter Schure schrijft over de Accumulation Bias. In ons laatste nummer kon u lezen hoe deze techniek gebruikt wordt om COVID-19-onderzoeken te combineren, nu gaat Judith gedetailleerd op de techniek in.

'Hardlopers zijn doodlopers' is een uitdrukking die aangeeft dat men niet te hard van stapel moet gaan, maar de krachten moet sparen om het tot het eind toe vol te houden. Bryan van Ingen heeft de gegevens van bijna 50.000 marathonlopers geanalyseerd om te kijken welke strategieën hiervoor bestaan.

Onze columnisten hebben niet stil gezeten: John Poppelaars schrijft over het proces van het modelleren, Henk Tijms ontzenuwt verdenkingen van vals spelen en Gerrit Stermerdink stelt voor om te bezuinigen op de Amerikaanse verkiezingen. Verder nieuws over de aan Ronald Does toegekende Lancaster Medal en een interview met Joaquim Gromicho die is benoemd tot hoogleraar aan de UvA.

Ook bevat dit nummer alle informatie over de virtuele Annual Meeting van de VVSOR en wordt de nieuwe voorzitter Casper Albers aan ons voorgesteld. De redactie bedankt scheidend voorzitter Fred van Eeuwijk voor zijn enthousiaste steun voor ons blad.

Wij besluiten met de inmiddels traditionele woorden 'wij wensen u veel leesplezier'.



INHOUD

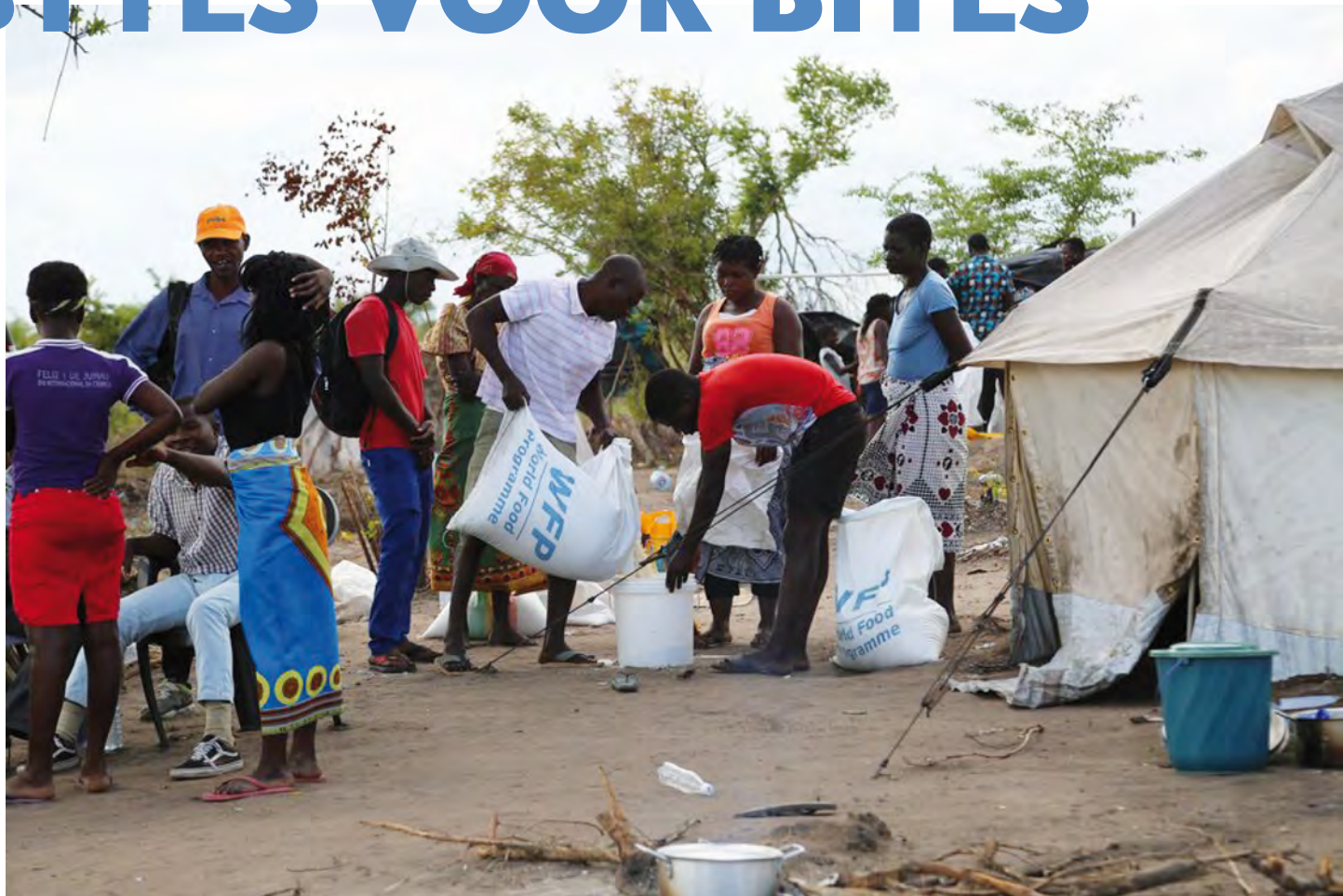
- 2 Ons vak heeft impact
- 4 Bytes voor bites; Naar een wereld zonder honger | MARLEEN BALVERT, MIRIAM CROUSEN, HEIN FLEUREN, PERRY HEIJNE & MELISSA KOENEN
- 9 Vaccinatielogistiek essentieel om vaccinverspilling en vertragingen te beperken | JAN FRANSOO & NIELS AGATZ
- 11 Ronald Does wint Lancaster Medal voor zijn internationale bijdragen aan kwaliteitsverbetering
- 12 Met een beetje mathematische besliskunde kom je een heel eind | JOAQUIM GROMICHO
- 15 Letter from the president elect

Going Viral VVSOR Annual Meeting 2021
Modeling spread Going Viral; Modeling Spread

- 16 Program Annual Meeting & Conference on March 18, 2021 | ONLINE

- 20 Bezuinigen op verkiezingen met slimme wiskunde – column | GERRIT STERMERDINK
- 21 Accumulation Bias; How to handle it ALL-IN | JUDITH TER SCHURE
- 28 Radicaal – column | JOHN POPPELAARS
- 30 Hoe loop je een marathon? | BRYAN VAN INGEN
- 34 Het balletje kan raar rollen – column | HENK TIJMS
- 35 Met analytics de wereld verbeteren

BYTES VOOR BITES



Mozambique. Foto: WFP/Rafael Tarasantchi

NAAR EEN WERELD ZONDER HONGER

Het hongerprobleem in de wereld is groot. Volgens de Verenigde Naties (VN) gaan maar liefst 690 miljoen mensen met honger naar bed. Elke tien seconden sterft ergens in onze wereld een kind als gevolg van honger en ondervoeding. Het Zero Hunger Lab van Tilburg University helpt om wereldwijde voedselzekerheid te realiseren door het gebruik van wiskundige technieken. Wij noemen dat 'bytes for bites'. Samen met anderen heeft Zero Hunger Lab inmiddels een twintigtal projecten op kunnen starten; een aantal leidde al tot meetbaar succes.

MARLEEN BALVERT, MIRIAM CROUSEN, HEIN FLEUREN, PERRY HEIJNE & MELISSA KOENEN

Om voor iedereen een betere en duurzamere toekomst te creëren, zijn er wereldwijd afspraken gemaakt en vastgelegd in de 17 zogeheten Sustainable Development Goals van de Verenigde Naties. Het tweede doel van de SDG's, SDG-2, is Zero Hunger met de volgende afspraken:

- 2.1 In 2030: veilig, voedzaam en voldoende voedsel voor iedereen
- 2.2 In 2030: een einde aan ondervoeding
- 2.3 In 2030: het verdubbelen van de landbouwproductiviteit en het inkomen van kleinschalige voedselproducenten
- 2.4 In 2030: duurzame voedselsystemen van boer-tot-bord
- 2.5 In 2030: instandhouding van de biodiversiteit van zaden, planten en dieren

Onze voedselsystemen moeten het komende decennium fundamenteel worden herzien, niet alleen om mensen in nood beter te kunnen helpen, maar ook om in 2050 op een duurzamere manier meer dan 10 miljard mensen dagelijks gezonde maaltijden te kunnen bieden.

Data science voor voedselzekerheid

Zero Hunger Lab van Tilburg University draagt met behulp van data science bij aan het realiseren van de SDG-2 afspraken. Dit doen we door samen te werken met hulporganisaties, ontwikkelingsorganisaties, bedrijven, overheid en kennisinstellingen die zich ook inzetten voor een wereld zonder honger. We helpen hen door het gebruik van data-analyses en het ontwikkelen van nieuwe algoritmen om betere beslissingen te nemen in een steeds meer complexe en onzekere wereld. Met andere woorden: hoe kun je de kracht van data science inzetten voor een betere wereld en een gezonde planeet?

Onze missie is mensen onafhankelijk te maken van voedselhulp zodat ze zelf kunnen zorgen voor duurzame voedselzekerheid. Dat doen we niet alleen in Afrika, Azië of het Midden-Oosten, maar ook in Nederland waar meer dan 150.000 mensen voor hun 'dagelijks brood' afhankelijk zijn van voedselbanken. In het nu volgende laten we vier projecten zien waar we mee bezig zijn.

Model OPTIMUS: oplossing voor efficiëntere voedselhulp

Het Wereldvoedselprogramma (WFP) van de VN helpt wereldwijd mensen in crisissituaties. Sinds 2011 ondersteunt Hein Fleuren (hoogleraar Operations Research aan de Tilburg University) met zijn studenten de VN bij het optimaliseren van hun *supply chain* in deze crisissituaties. De intensieve samenwerking resulteerde uiteindelijk in het door Koen Peters ontwikkelde OPTIMUS, een innovatieve oplossing voor effectievere voedselhulp, die door het WFP gebruikt wordt in landen als Jemen, Syrië en Zuid-Soedan.¹

Het model achter OPTIMUS is een geheeltallig lineair programmeringsprobleem met als belangrijkste basis een netwerkmodel. Het neemt alle aspecten mee die relevant zijn voor de voedselketen van het WFP. Dit houdt in de manier waarop de begunstigden worden voorzien van voedsel (voedsel- of geldhulp), de inhoud van een voedselpakket zodat de begunstigden een gezond en betaalbaar dieet ontvangen, welke producent de hulp levert en de manier en routing van transport van de gekozen hulpvorm. Door onder andere de kosten te minimaliseren van het leveren van hulp, kunnen meer mensen met het beschikbare budget worden geholpen.

- De concrete resultaten van OPTIMUS op een rijtje:
- Er kunnen tot wel 15–20% meer mensen gevoed worden met hetzelfde budget.
 - Er is al meer dan 30 miljoen dollar bespaard door het gebruik van OPTIMUS.
 - Na gebruik in tal van landen wil het WFP het model in 80 landen inzetten.

Child Growth Monitor: adequaat meten van ondervoeding

Ongeveer 200 miljoen kinderen gaan met honger naar bed volgens statistieken van FAO en de Wereldbank. Het tijdig opsporen van ondervoeding bij kinderen door middel van adequate metingen is cruciaal voor de behandeling. Uit de statistieken van de Wereldgezondheidsorganisatie blijkt dat slechts 35% van de ondervoede kinderen adequaat wordt gemeten, waardoor hulpgoederen niet optimaal worden ingezet.

Welthungerhilfe ontwikkelt met onder andere Zero Hunger Lab een app die, op basis van een foto of filmpje gemaakt met de mobiele telefoon, ondervoeding detecteert. Hierdoor kan snel, veilig en overal ter wereld het niveau van ondervoeding van kinderen worden vastgesteld, en kan voedselhulp sneller en beter worden verleend.

Tilburg University's Zero Hunger Lab ondersteunt Welthungerhilfe bij de ontwikkeling van het diepe neurale netwerk dat op basis van beeldmateriaal van een kind voorspelt of het ondervoed is. Voor Welthungerhilfe is het belangrijk dat de app voor zo veel mogelijk mensen eenvoudig te gebruiken is. Een belangrijk onderdeel daarvan is de vraag aan welke eisen een foto moet voldoen zodat een neurale netwerk een goede voorspelling kan maken. Hierbij kan gedacht worden aan eisen voor de achtergrond, de afstand van de camera tot het kind, vanuit welke richting het kind wordt gefotografeerd, en nog veel meer. Hoe korter de lijst met eisen is, hoe toegankelijker de app. Zero Hunger Lab onderzoekt welke randvoorwaarden werkelijk nodig zijn, en welke losgelaten kunnen worden. Zo wordt de Child Growth Monitor toegankelijk voor een groot aantal mensen en kunnen kinderen met ondervoeding sneller geïdentificeerd worden.



ENHANCE: healthy people on a healthy planet

Onze huidige voedselsystemen zijn er niet op ingericht om de groeiende wereldbevolking te voorzien van gezonde diëten. Er is een disbalans, waarbij er enerzijds mensen overvoed zijn en een derde van ons voedsel wordt weggegooid en anderzijds mensen ondervoed zijn. Daarnaast wordt er nog te weinig aandacht besteed aan de mate waarin voedselsystemen verantwoordelijk zijn voor klimaatproblemen, aantasting van ecosystemen en biodiversiteit.

Samen met Johns Hopkins University, Capgemini en WFP werken we aan een optimalisatiemodel om met name overheden te kunnen adviseren. Dit model is gebaseerd op het dieetprobleem: één van de klassieke voorbeelden van lineair programmeren. Het model moet niet alleen rekening houden met gezonde eetpatronen, culturele voorkeuren en de kosten van het dieet, maar ook met het minimaliseren van de impact op het milieu.

De eerste pilotmodellen worden in 2020/2021 ontwik-

keld voor Indonesië. Hierbij richt het Zero Hunger Lab zich met name op twee onderdelen. Het eerste is duurzaamheid. Het huidige model voor het optimaliseren van diëten minimaliseert de kosten zodanig dat aan nutritionele eisen is voldaan en het dieet aan de culturele wensen voldoet. Dit model wordt uitgebreid met een duurzaamheidsdoelstelling: het samengestelde dieet moet rekening houden met de invloed op het klimaat (broeikasgassen), watergebruik en bodemgesteldheid. Dit resulteert in een multi-criteria optimalisatieprobleem. Daarnaast wordt prijsrobuustheid in het model verwerkt: zelfs bij grote veranderingen in bijvoorbeeld prijs en beschikbaarheid van voedsel moeten mensen met weinig of geen inkomen toegang blijven houden tot gezond voedsel.

Met andere woorden: 'healthy people on a healthy planet'. In de komende jaren wil de unieke alliantie de innovatieve ENHANCE-oplossingen verder ontwikkelen en implementeren voor landen in Afrika. Hiervoor wordt actief gezocht naar strategische partners.

Voedselbanken Nederland

In Nederland voorzien 171 voedselbanken zo'n 150.000 cliënten van voedsel. Het Zero Hunger Lab helpt Voedselbanken Nederland meer informatie te halen uit hun eigen en publieke data, zodat de hulpvoorziening verder verbeterd kan worden. Zo zijn er naar verwachting zo'n 200.000 tot 250.000 mensen die recht hebben op hulp van de voedselbank, maar hun weg daarheen nog niet hebben weten te vinden. Het Zero Hunger Lab helpt Voedselbanken Nederland met behulp van *machine learning* en publieke data inzicht te krijgen in wie deze mensen zijn, welke eigenschappen ze karakteriseren en in welke regio's ze wonen. Daarnaast gebruiken we machine learning om Voedselbanken Nederland inzicht te geven in hoe ze hun website kunnen verbeteren om de voedselbank toegankelijker te maken. Iedere voedselbank heeft haar eigen website. Een masterstudent van Tilburg University heeft kenmerkende eigenschappen van deze websites in kaart gebracht, en geanalyseerd welke kenmerken ertoe leiden dat mensen eenvoudig contact kunnen opnemen met hun lokale voedselbank.

Ook helpt het Zero Hunger Lab de voedselbanken bij

DOE MEE

Iedereen kan meehelpen. Het Zero Hunger Lab zoekt continu gedreven masterstudenten, PhD's en specialisten die zich willen inzetten voor Zero Hunger.

- Als vrijwilliger kun je meehelpen bij bepaalde analyses. Dikwijls hebben we voor grote projecten een voor-analyse nodig.
- Als specialist kun je meehelpen door mee te denken met specifieke vraagstukken gerelateerd aan methoden.
- Als partner kun je ons helpen door ons te betrekken bij gefinancierde wetenschappelijke onderzoeken (NWO, EU, etc); of wellicht weet je generieke fondsen die voor ons van belang kunnen zijn.

Let wel dat we vaak niet direct geschikte projecten of vragen hebben. We zetten onze 'vrienden van Zero Hunger Lab' op een lijst en bekijken heel regelmatig wat we nodig hebben afhankelijk van de vragen van onze partners.

Het Zero Hunger Lab is in 2019 opgericht naar aanleiding van de impactvolle samenwerking met het World Food Programme (WFP). Met allerlei vormen van data science willen we bijdragen aan het reduceren van honger; niet alleen door noodhulp te verbeteren maar ook door mee te analyseren en optimaliseren hoe voedsel beter verbouwd, bewaard en niet verspild kan worden. Kijk op www.tilburguniversity.edu/nl/zerohungerlab als je meer wilt weten en hoe je samen met ons kunt bijdragen aan het oplossen van het wereldhongerprobleem

Voor vragen neem contact op met: prof.dr.ir. Hein Fleuren, fleuren@tilburguniversity.edu.

het verbeteren van hun *supply chain*. De voedselbanken hebben te maken met grote onzekerheden: aan de vraagkant fluctueert het aantal cliënten, wat versterkt is door de huidige coronapandemie, en aan de aanbodkant varieert het aantal beschikbare vrijwilligers en de hoeveelheid en het soort voedsel dat gedoneerd wordt. We verminderen de onzekerheid door met behulp van data-analyse inzicht te geven in het aantal extra cliënten waar de voedselbanken op moeten rekenen naar aanleiding van corona; en we werken aan robuuste optimalisatiemodellen om de supply chain te versterken. Daarnaast kijken we met behulp van supply-chainmodellen en optimalisatie naar de mogelijkheden van het herverdelen van voedingsmiddelen tussen de lokale voedselbanken. Kortom, een grote variatie aan technieken helpt Voedselbanken Nederland haar hulpverlening te versterken.

We doen het samen

Om het hongerprobleem in de wereld te verkleinen is niet alleen efficiëntere en effectievere (nood)voedselhulp noodzakelijk. Belangrijk is het de lokale capaciteit te versterken zodat boeren, bedrijven en gemeenschappen zelf voor duurzame voedselzekerheid kunnen zorgen en ze onafhankelijk van steun worden. Het is daarom nodig samen te werken en kennis en ideeën te delen.

Het Zero Hunger Lab werkt met een eigen team zelfstandig aan onderzoek en oplossingen voor het wereldhongerprobleem. Binnen het lab zijn meer dan twintig studenten en onderzoekers van Tilburg University actief. Naast Tilburg University heeft het ministerie van Buitenlandse Zaken zich als strategisch partner verbonden aan het meerjarige researchprogramma van het Zero Hunger Lab.²

De oplossing van het voedselprobleem ligt in samenwerking en gezamenlijk oplossingen ontwikkelen (co-creatie). Voor ons Lab zijn twee dingen belangrijk:

1. Bereiken we met onze modellen, analyses en algoritmen werkelijk een verbetering in het terugdringen van honger?
2. Doen we dit op een wetenschappelijk goede en verantwoorde manier?

Het Zero Hunger Lab is nadrukkelijk geen plaats waar we vooral consultancy of korte opdrachten doen, maar het is ook geen plaats waar we alleen streven naar 'toppublicaties'. We richten ons vooral op langetermijnpact door goed onderzoek.

Wat hebben we al bereikt?

Door co-creatie met andere kennisinstellingen, hulp- en ontwikkelingsorganisaties hebben we inmiddels meer dan twintig onderzoeksprojecten op kunnen starten. Een aantal projecten leidt al tot concrete impact zoals het eerdergenoemde project OPTIMUS. Koen Peters, PhD-researcher van Zero Hunger Lab leidt diverse initiatieven voor WFP en is met het werk voor OPTIMUS één van de finalisten van de INFORMS Franz Edelman Award. De prijs is in het leven geroepen om erkenning te verlenen aan belangrijk onderzoek op het gebied van 'analytics and operations'.

Andere hebben, ondanks dat ze nog niet zijn geïmplementeerd, enorme potentie zoals de Child Growth Monitor. Weer andere zijn in de conceptfase zoals ENHANCE en Onder de Radar 2.0. Ook deze kunnen een belangrijke bijdrage leveren aan een wereld zonder honger.

1. Koen Peters, PhD-researcher van ZHL, leidt diverse initiatieven voor WFP en is met het werk voor OPTIMUS één van de finalisten van de Franz Edelman Award.

2. Zero Hunger Lab werkt samen met tal van hulp- en ontwikkelingsorganisaties en kennisinstellingen zoals: World Food Programme, Solidaridad, Welthungerhilfe, Voedselbanken Nederland, Oxfam, Dutch Relief Alliance, Landelijk Operationeel Team-Corona, Wereld Bank, One-Acre Fund, INSEAD Humanitarian Research Group, Wageningen University & Research, Dutch Coalition for Humanitarian Innovation, Center for Frugal Innovation Africa, KU Leuven, University of Liberia.

MARLEEN BALVERT is assistant professor aan Tilburg University en het Zero Hunger Lab.
E-mail: m.balvert@tilburguniversity.edu

MIRIAM CROUSEN is Communications Office Manager bij Zero Hunger Lab.
E-mail: m.crousen@tilburguniversity.edu

HEIN FLEUREN is scientific director van het Zero Hunger Lab en hoogleraar Toepassen van Operations Research.
E-mail: fleuren@tilburguniversity.edu

PERRY HEIJNE is co-founder van Zero Hunger Lab.
E-mail: P.C.Heijne@tilburguniversity.edu

MELISSA KOENEN is PhD student bij het Zero Hunger Lab van Tilburg University.
E-mail: m.f.koenen@tilburguniversity.edu

Vaccinatielogistiek essentieel om vaccinverspilling en vertragingen te beperken



Met het oog op de opmars van de Britse variant van het coronavirus en daarmee de dreiging van een derde golf, is het nog meer zaak om de logistieke vaccinatie-operatie zo in te richten dat er weinig of geen vaccins worden verspild

JAN FRANSOO & NIELS AGATZ

Na een buitengewoon snel proces van ontwikkelen is het eerste Coronavaccin op 21 december 2020 goedgekeurd in Europa. De volgende stap is het snel, veilig, en zonder onnodig vaccinverspilling vaccineren van miljoenen Nederlanders. Dat is een ongekennde uitdaging vanwege de grote onzekerheid in vraag en aanbod, de enor-

me schaal en de complexe logistieke randvoorwaarden. Dit laatste heeft onder andere betrekking op de beperkte houdbaarheid en benodigde koeling. De aantallen vaccins die Nederland verwacht te ontvangen zouden ons in staat moeten stellen om voor de zomer van 2021 het niveau van groepsimmuniteit te hebben bereikt. Na de

unieke snelheid waarmee de vaccins ontwikkeld zijn, mag het eenvoudigweg niet zo zijn dat er door logistieke problemen een minder dan perfecte vaccinatiecampagne plaatsvindt. Dat is ook niet nodig. Fundamentele en reeds beproefde inzichten uit de logistieke wetenschap kunnen bijdragen aan een snelle en doeltreffende uitrol van de vaccinaties én het minimaliseren van vaccinverspilling. Het door het RIVM in de Kamerbriefing van 9 december 2020 verwachte vaccinverspilling van 10 tot 20% (onder de term ‘spillage’) is in onze ogen niet alleen onaanvaardbaar, maar ook onnodig. De volgende inzichten zijn daarbij van belang.

Zo veel mogelijk werken met grote locaties leidt tot zo min mogelijk vaccinverspilling

Bij grote onzekerheid over vraag en aanbod kan de flexibiliteit vergroot worden door logistieke operaties te centraliseren. Dat betekent dat het beter is om op minder locaties te vaccineren en voorraad te houden. Binnen het Coronavaccinatieprogramma hebben we te maken met veel onzekerheid, zowel qua vraag als aanbod. Het is moeilijk in te schatten hoeveel en welke mensen gevaccineerd zouden willen worden. Door op een beperkt aantal van 30 tot 50 centrale locaties voorraad te beheeren hoeft alleen maar een totaalinschatting te worden gemaakt voor enkele honderdduizenden mensen per locatie. Indien dit bij duizenden individuele huisartsenposten moet gebeuren, is de relatieve variatie en onzekerheid veel groter, en zouden overschotten bij een huisarts (door lage vaccinatiebereidheid in haar praktijk bijvoorbeeld) moeten worden opgehaald en versleept naar een andere huisarts die wellicht te maken heeft met een hoge vaccinatiebereidheid. De stelling van het ministerie van VWS dat Nederland klein genoeg is om dit te regelen, is hierbij irrelevant: heen en weer verslepen van kleine overschotten aan vaccins tussen duizenden huisartsen zal tot onnodig veel verlies aan vaccin leiden: het kost al snel een dag en dat is veel als de houdbaarheid maar enkele dagen is. Bovendien leidt dit tot extra overdrachtpunten, wat de kansen op beschadiging en het kort verbreken van de temperatuureisen vergroot. Als er ruim voldoende vaccins zouden zijn – zoals in

de meeste jaren bij de influenzacampagne – zijn die overschotten geen probleem. We zouden ons kunnen permitteren om kleine hoeveelheden te moeten vernietigen omdat de houdbaarheid is verstreken. Nu we bij de COVID-19-vaccinatie nog tenminste een half jaar met tekorten geconfronteerd zullen worden, is onnodig verlies van vaccins onwenselijk

Elke decentrale vaccinatielocatie koppelen aan een van de grote vaccinatiecentra vermindert de kans op vaccinverspilling

Het beheren van de voorraad op centrale locaties betekent niet automatisch dat we niet kleinschalig kunnen vaccineren waar het echt niet anders kan. Dat is vooral van belang voor diegenen die onmogelijk naar een centrale locatie kunnen komen zoals mensen met een zwakke gezondheid. Kleine locaties zoals verpleeghuizen, instellingen en huisartsen kunnen door frequente leveringen voorzien worden van vaccins vanuit de grotere locaties. Logistiek robuust is het om deze fijnmazige distributie te regelen vanuit dezelfde 30-50 regionale grote locaties zodat de voorraad makkelijk kan worden verdeeld bij overschotten op de decentrale locaties zoals een verpleeghuis. Een dergelijk hub-and-spokesysteem is heel gebruikelijk in de logistiek (bijvoorbeeld bij het beleveren van supermarkten) en er is veel expertise in Nederland om dit goed in te richten. Hierbij moeten we ons wel realiseren dat we moeten proberen om het aantal decentrale locaties en de daarbij behorende logistiek tot een minimum te beperken. Het ligt daarbij voor de hand om bijvoorbeeld wel de 700 verpleeghuizen mee te nemen met al hun bewoners; gegroepeerd rondom 30 centrale locaties gaat het dan om ongeveer 25 verpleeghuizen per ‘hub’ – opnieuw een behapbare complexiteit.

Grotere vaccinatiecentra maken het systeem flexibel en robuust

Bij grote onzekerheid is het noodzakelijk om op tijd voorbereiding te treffen. Afwachten tot alle details dui-

delijk zijn betekent vaak dat we vervolgens niet snel kunnen schakelen. Bij het inschalen van de capaciteit is het verstandig om rekening te houden met zo veel mogelijk scenario's. We moeten daarbij kijken naar zowel de kansen en consequenties van verschillende capaciteitskeuzes in verschillende scenario's. Te veel capaciteit betekent meer mensen inroosteren en ruimtes inplannen dan wat strikt noodzakelijk zou zijn in een voorspelbaar en stabiel systeem. Extra capaciteit geeft echter ook de flexibiliteit om snel veel mensen te kunnen vaccineren en beter om te gaan met schokken in vraag en aanbod. Te weinig capaciteit kan er voor zorgen dat we vertraging oplopen bij het vaccineren. Te weinig capaciteit kan mogelijk ook leiden tot lange wachttijden op de locaties met hoger besmettingsrisico. Geven de grote maatschappelijke en economische schade die onze samenleving dag na dag oploopt, en de relatief lage kosten van overcapaciteit, is het verstandig om te plannen op forse overcapaciteit.

Ten slotte maken wij ons bij een verregaande decentrale strategie zorgen over de sterk variërende belasting bij de huisartsenpraktijken. De vele onzekerheden bij het coronavaccin zullen ertoe leiden dat een huisarts niet zoals bij de griep prik enkele keren gepland een aantal zendingen zal ontvangen en haar cliënten in korte tijd kan inenten. Het is zeer waarschijnlijk dat de levering in grote onvoorspelbare schokken zal komen en steeds op zeer korte termijn relatief kleine groepen moeten worden ingeënt. Het zou voor iedereen beter zijn dat zoveel mogelijk te vermijden.

Dit artikel verscheen eerder in *Medisch Contact* en is door de auteur aangepast voor *STATOR*. Na het afronden ervan is bekend geworden dat een groep deskundigen onder leiding van Jan Fransoo het RIVM adviseert over de distributie, zie: <https://www.tilburguniversity.edu/current/news/more-news/jan-fransoo-advises-rivm>

JAN FRANSOO is hoogleraar Operations en Logistiek Management aan Tilburg University en aan de Technische Universiteit Eindhoven.
E-mail: Jan.Fransoo@tilburguniversity.edu

NIELS AGATZ is universitair hoofddocent Transport en Logistiek aan de Erasmus Universiteit.
E-mail: nagatz@rsm.nl

Ronald Does wint Lancaster Medal voor zijn internationale bijdragen aan kwaliteitsverbetering

Als eerste Nederlander ooit wint Ronald Does, hoogleraar Industriële Statistiek aan de Universiteit van Amsterdam, de Lancaster Medal. Deze prestigieuze medaille van de American Society for Quality (ASQ) wordt sinds 1981 jaarlijks uitgereikt aan mensen die zich internationaal zeer verdienstelijk hebben gemaakt om de kwaliteit van processen, producten en diensten te verbeteren. Eerder winaars zijn Armand Feigenbaum, grondlegger van Total Quality Management, en Noriaki Kano, bedenker van het klantverwachtingenmodel.

Does studeerde wiskunde in Leiden, waar hij later ook promoveerde. Na zijn promotie ging hij aan de slag als universitair docent, later als hoofddocent, medische statistiek en psychometrie bij de Universiteit Maastricht. Hij werkte daarna bij Philips aan het verbeteren van de kwaliteit van producten, door met behulp van statistische methodes kwetsbare plekken in het productieproces op te sporen.

Vanaf 1991 is Does hoogleraar aan de UvA. In 1994 richtte hij het Instituut voor Bedrijfs- en Industriële Statistiek (IBIS UvA) op, van waaruit hij bedrijven als ABN AMRO Bank, ASML, DAF Trucks, Douwe Egberts, Stork, Verkade, Zwitserleven en tientallen ziekenhuizen hielp om de bedrijfsprocessen te verbeteren. Daarnaast is Does grondlegger van de opleidingen Lean Six Sigma in Nederland. Hij zorgde ervoor dat duizenden mensen een opleiding tot Green Belt of Black Belt afrondden. Uniek aan de opleiding is dat deze een praktijkproject omvat, zodat het bedrijf de kosten van de opleiding van een medewerker ruimschoots terugverdient.

Does combineerde altijd praktijk en wetenschap. Hij schreef 125 wetenschappelijke artikelen en 10 boeken over industriële statistiek en begeleidde 24 promovendi, waaronder 3 promovendi uit Pakistan.

Internationaal was hij actief betrokken bij de oprichting van the International Society for Business and Industrial Statistics (ISBIS), the European Network for Business and Industrial Statistics (ENBIS), and the International Statistical Engineering Association (ISEA). Does is Fellow van de American Statistical Association en de American Society for Quality. Eerder ontving hij de William G. Hunter Award en de Walter A. Shewhart en George E.P. Box Medals.

In 2021 gaat Does met emeritaat. Zijn afscheidsrede zal gaan over zijn ervaringen als statisticus in de afgelopen 45 jaar.



Foto: Tom van Limpt

MET EEN BEETJE MATHEMATISCHE BESLISKUNDE KOM JE EEN HEEL EIND

JOAQUIM GROMICHO

We zitten nu – februari 2021 – in de tweede lockdown tijdens de voortslepende COVID-19-pandemie. Voor de tweede keer in minder dan een jaar – dat twee opeenvolgende schooljaren beslaat – zitten alle leerlingen de hele dag thuis, met uitzondering van kwetsbare kinderen en kinderen van ouders met vitale beroepen. Een grote beproeving voor de kinderen, de leerkrachten en uiteraard ook voor de ouders. Dat was het ook tijdens de eerste lockdown in het voorjaar van 2020. Toen het kabinet op

21 april 2020 besloot om de basisscholen vanaf 11 mei weer te openen, was iedereen blij, maar voor scholen was het een flinke (wiskundige) puzzel om alles te organiseren. Scholieren zouden ongeveer de helft van hun lestijd weer les op school gaan krijgen, maar wel in kleinere groepen. De andere helft van de tijd zouden ze thuis de lessen volgen.

Zelf heb ik drie dochters op dezelfde basisschool. In het schooljaar 2019–2020 zaten er 553 kinderen op die school,

Hoe een wiskundige de school van zijn kinderen tijdens de lockdown helpt bij maken van een indeling van de klassen waarbij rekening gehouden moet worden met de wensen van de ouders en de leerkrachten.

afkomstig uit 407 gezinnen en verdeeld over 22 groepen. Elke groep moest dus in twee delen worden gesplitst.

Voor elk kind moest het besluit worden genomen: deel A of deel B? Deel A ging naar school op maandag en donderdag en op de oneven weken op woensdag, deel B op dinsdag en vrijdag en op de woensdagen van de even weken.

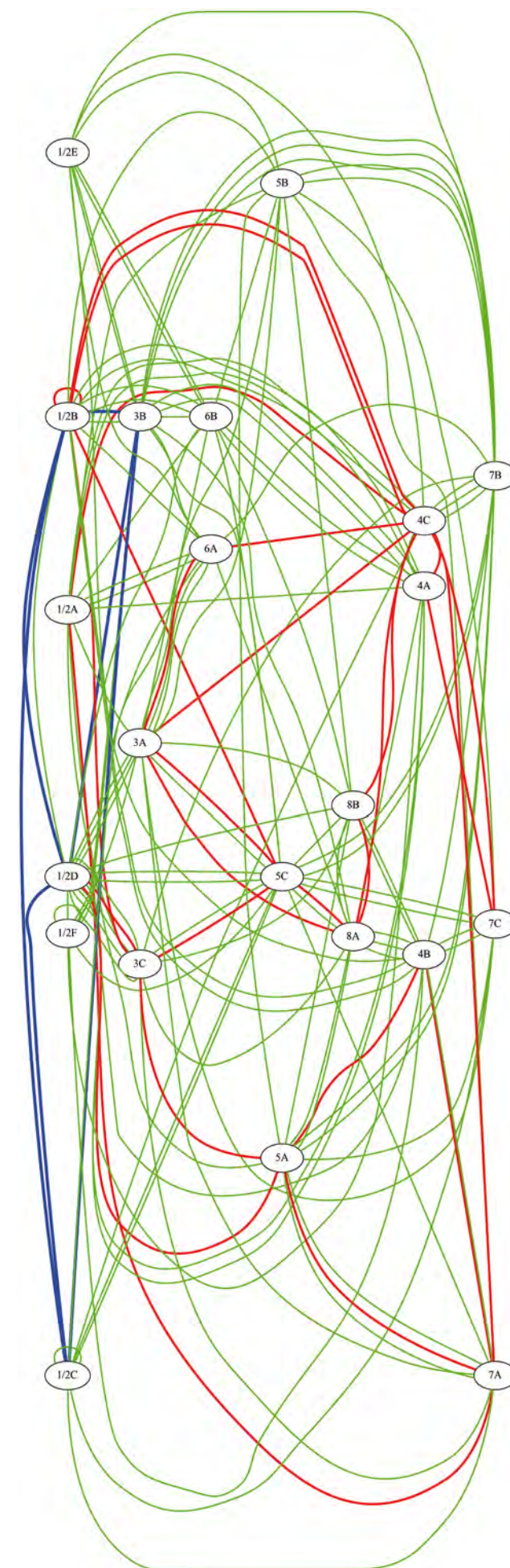
Het was een nadrukkelijke wens van de ouders dat alle kinderen uit hetzelfde gezin op dezelfde dagen naar school zouden gaan, oftewel in hetzelfde deel zou komen. Ook zouden de ouders met vitale beroepen mogen kiezen in welk deel hun kinderen zouden komen. Daarnaast was er de wens van de leerkrachten (en de kinderen zelf) dat de splitsing rekening zou houden met de sociale cohesie onder de kinderen. Ze hadden het al zwaar te verduren, dus dan maar het liefste samen met de hartsvriendinnen en de boezemvrienden. Verder was de wens dat voor groepen met meerdere deeltijdleerkrachten het aantal verschillende leerkrachten per groep per week zo laag mogelijk zou zijn. Uiteraard moesten ook de deelgroepen van gelijke grootte zijn (plus minus een voor oneven aantallen). En laatst maar niet het minst, gesteld dat een leerkracht van risicogroepen vervangen moest worden, dat er ook dan een oplossing met evenveel lesdagen op school voor alle kinderen zou worden gevonden, idealiter zou ieder kind een vol schema A of B krijgen.

Alleen al de wens ‘alle kinderen van hetzelfde gezin op A of op B’ was voor menig schooldirecteur onbegonnen werk om met pen en papier op te lossen. Neem bijvoorbeeld de graaf die voor de groepen van de school van mijn kinderen aangeeft met een groene lijn dat twee kinderen van hetzelfde gezin in de verbonden groepen zitten of met een rode lijn dat er alle drie kinderen van hetzelfde gezin zich over die groepen verdelen en in blauw voor de vier kinderen uit hetzelfde gezin. Dan begrijpen we dat alleen al deze doelstelling de pet van de gemiddelde overbelaste schooldirecteur ver te boven gaat.

Gesneden koek

Maar zo’n puzzel is gesneden koek voor wie een beetje mathematische besliskunde heeft geleerd. Het meest bekende deel ervan, gemengde geheeltallige lineaire optimalisatie, kan deze puzzel namelijk prima oplossen.

Gewapend met een Excelbestand van gezinssamenstellingen, met Python en met mijn kennis van besliskunde kon ik aan alle wensen voldoen en zonder moeite de puzzel meerdere malen opnieuw oplossen, telkens na het ontstaan van nieuwe uit voortschrijdend inzicht geboren wensen.*



Mijn model implementeerde alle wensen als zachte restricties, dit om een toegelaten oplossing te kunnen garanderen. De gewichten waarmee schendingen in de doelfunctie werden beboet waren gekozen overeenkomstig de prioriteit van de wensen. De allerbelangrijkste wens was dat alle kinderen van hetzelfde gezin op dezelfde dagen naar school zouden gaan.

Als belangrijkste beslissingsvariabelen heb ik x_{ij} voor elke leerling i en elk deel $j \in \{A, B\}$ binair genomen: voor elke kind i moet idealiter een van x_{iA} of x_{iB} gelijk aan 1 zijn, en de andere is 0. Om een zachte restrictie te implementeren werd $x_{iA} + x_{iB} \leq 1$ genomen, terwijl de som over alle kinderen werd gemaximaliseerd. Voor elke gezin G worden voor alle kinderen $i \in G$ alle x_{iA} gelijk aan elkaar gesteld en ook x_{iB} . Om de delen zo gelijk aan elkaar mogelijk te krijgen werden twee variabelen per groep gedefinieerd voor de aantal kinderen in de kleinste en de grootste deel respectievelijk en het totale verschil werd geminimaliseerd. Alle andere randvoorwaarden werden op een gelijke manier gemodelleerd.

Dit is de gesimplificeerde code (zie ook de hierbij afgebeelde code) waarbij elk gezin in een deel als harde restrictie wordt gemoduleerd en enkel naar delen van zo gelijk mogelijke grootte wordt geoptimaliseerd. Het hierboven geschetste probleem is binnen 1,5 seconden op te lossen met de gratis CBC-solver (<https://github.com/coin-or/Cbc>).

De volledige puzzel – tegemoetkomend aan **alle wensen** wordt met dezelfde gratis *solver* in minder dan een minuut opgelost. Om de veranderingen niet te groot te

laten zijn werd, bij elke toevoeging van een nieuwe wens, getracht om de nieuwe oplossing zo gelijk mogelijk aan de vorige versie te krijgen, dat wederom met zachte restricties. Als extra service heb ik, omdat ik toch in de Pythoncode zat, automatisch door middel van PyLaTeX de indelingsbrieven in pdf laten genereren. Geen handwerk, geen knip- en plakwerk, geen fout.

Toen Thomas Davenport in 2012 data scientist het meest sexy beroep op aarde noemde, presenteerden zich in korte tijd op LinkedIn enorme aantallen data scientists. Helaas beschikt slechts een minderheid daarvan over serieuze kennis van wiskunde, statistiek of (nog minder!) mathematische besliskunde. Dat is jammer, want alleen dan kan men dit soort problemen uit het dagelijks leven op een snelle en gemakkelijke manier oplossen.

In deze tijden van lockdown raad ik van harte het geweldige boek *Operationele analyse* (2002) van onze trouwe columnist Henk Tijms aan. Samen met *Pyomo; optimization modeling in Python* van William Hart en Carl D. Laird (2017) komt u een heel eind!

* De medezeggenschapsraad – waarvan ik deel uitmaak – gaf toestemming om de gegevens over gezinssamenstelling onder strikte voorwaarden en enkel voor dit doel te gebruiken. De gegevens zijn anoniem en onherkenbaar verwerkt.

JOAQUIM GROMICHO is sinds januari 2021 hoogleraar Business Analytics aan de Universiteit van Amsterdam. Ook is hij Science & Education Officer bij ORTEC en hoofdredacteur van *STATOR*. E-mail: Joaquim.Gromicho@ortec.com



VVSOR shows a growing collaboration between disciplines focused on the visibility of statistics and operations research. Despite the troubling times we all faced during 2020, there was a huge increase in the presence of data analysis in society. Never before were phrases like “curve flattening”, “false positive diagnostic tests” and “vaccine efficacy” so widespread.

Several of our members became the face of the scientific response to the pandemic, of the cross-discipline collaboration, data literacy and healthy skepticism. One noteworthy example is prof. dr. Casper Albers. With his columns for the newspaper *de Volkskrant*, Casper Albers encourages his readers to arm themselves with statistics to fight the threat of fake news. In the University of Groningen university newspaper, he empowers both students and employees to be the future of the university. And many news outlets seek his advice, like Nu.nl and *Het Dagblad van het Noorden*, or television shows like Human’s *Medialogica*. Casper is an important voice in the scientific landscape that has argued for critical changes to decrease sloppy research and for a bottom-up approach to scientific progress.

It is with great pride that we announce the candidacy of Casper Albers to be the next president of the VVSOR. At the Annual Meeting of March 18th, if the meeting votes in favour, Casper will succeed Fred van Eeuwijk who served as the VVSOR’s president for four years. Next, Casper gives us a quick introduction.

THE VVSOR DAILY BOARD



Foto: Bob Bronshoff

Letter from the president elect

During the Annual Meeting, the Society can elect me as successor of Fred van Eeuwijk. As the next President, I hope to contribute to the visibility of our society, towards our own members as well as towards society at large. The current daily board made great improvements to the administrative structure of our society. The new board can therefore focus on the future: how do we achieve a steady stream of new, young members and how do we stay relevant for the next ten years or so? I hope to contribute to those challenges.

CASPER ALBERS

Casper Albers received his PhD in Groningen in 2003, under the supervision of Willem Schaafsma. After two postdocs he returned to Groningen in 2009, to join the Faculty of Social and Behavioural Sciences. Here, he holds the chair on Applied Statistics and Data Visualisation. His research focuses on longitudinal models for psychological data and statistics communication. Within VVSOR, Casper has served as the president of the social science section for several years. Outside the VVSOR he fulfilled roles as Research Director of the Heymans Institute for Psychological Research, board member of the International Association for Statistical Computing and as a board member of the Groningen University Fund.

```
# example data
groups = { 1 : ['a', 'c'], 2 : ['b', 'd'] }
families = [ ['a', 'b'], ['c'], ['d'] ]

import pyomo.environ as pyo

m = pyo.ConcreteModel( 'split' )
m.part = [ 'A', 'B' ]
m.child = {c for f in families for c in f}
m.group = groups.keys()

def familyPairs(m):
    return [ (i,j) for f in families for i,j in zip(f[:-1],f[1:]) if len(f) > 1 ]
m.together = pyo.Set( dimen = 2, initialize = familyPairs )

m.x = pyo.Var( m.child, m.part, within=pyo.Binary )
m.minPart = pyo.Var( m.group, within=pyo.NonNegativeReals )
m.maxPart = pyo.Var( m.group, within=pyo.NonNegativeReals )

m.objective = pyo.Objective( expr = sum( m.maxPart[g]-m.minPart[g] for g in m.group ),
                             sense= pyo.minimize )

m.choose = pyo.Constraint( m.child,
                           rule = lambda m,c : sum( m.x[c,p] for p in m.part ) == 1 )
m.join = pyo.Constraint( m.together, m.part,
                         rule = lambda m,i,j,p : m.x[i,p] == m.x[j,p] )
m.below = pyo.Constraint( m.group, m.part,
                          rule = lambda m,g,p : sum( m.x[c,p] for c in groups[g] ) >= m.minPart[g] )
m.above = pyo.Constraint( m.group, m.part,
                          rule = lambda m,g,p : sum( m.x[c,p] for c in groups[g] ) <= m.maxPart[g] )
```

Code (gesimplificeerd) voor het indelen van schoolklassen met een aantal restricties

Going Viral

Modeling spread

Annual Meeting of the Netherlands Society for Statistics and Operations Research (VVSOR)

ONLINE

Thursday March 18, 2021

12:00 – 17:15

Four international speakers will reflect on modeling Going Viral and modeling spread, or how to reduce the spread of COVID-19, from their respective fields by using concepts from both statistics and OR. Invited speakers:

- ◆ Prof. dr. N.V. (Nelly) Litvak
- ◆ Prof. dr. G. (Gavin) Shaddick
- ◆ Prof. dr. D.J. (Dimitris) Bertsimas
- ◆ Prof. dr. J. (Jacco) Wallinga

For the first time, this year the Annual Meeting will be an online streaming event on Zoom and Youtube. We will have a general assembly for members, followed by the actual event with four talks and two award presentations. The AM 2021 will held in English.

Registration is open, please register on the vvsor-website <https://www.vvsor.nl/articles/vvsor-annual-meeting-2021>.

This year's annual meeting is free of charge, and members can receive a small snack box if they register early.

E-mail: annualmeeting@vvsor.nl

DATE

Thursday, March 18, 2021

VENUE

Online via Zoom and Youtube

REGISTRATION

Registration for the conference is mandatory at www.vvsor.nl/articles/vvsor-annual-meeting. Detailed information can be found on our website.

LANGUAGE

The talks at the annual meeting will be in English.

ALGEMENE LEDENVERGADERING (ALV)

The Annual General Meeting of members (ALV) takes place on March 18, 12:00 – 13:00. The relevant documents will be provided on the website two weeks before the meeting.

LUNCH, COFFEE, TEA AND DRINKS

The first 100 members who register before March 1st 2021 can ask for a snack box to be sent to their home address before the meeting.

ORGANIZING COMMITTEE

The annual meeting is organized by a special committee in cooperation with the board of the VVSOR. For questions, contact the administration at annualmeeting@vvsor.nl.

REGISTRATION ENDS AT MARCH 16!

12:00 – 13:00 **Annual General Meeting (ALV)**

13:00 – 13:30 Lunch break

13:30 – 13:45 **Welcome and Opening Talk** by CASPER ALBERS, President of the VVSOR

13:45 – 14:15 **Containment strategies for COVID-19 using mobility data**
NELLY LITVAK, *University of Twente* | *Eindhoven University of Technology*

14:15 – 14:45 **Incorporating time-use and health data into a Dynamic Microsimulation Epidemiological model for COVID-19**
GAVIN SHADDICK, *University of Exeter*

14:45 – 15:15 Break

15:15 – 15:45 **From predictions to prescriptions: A data-driven response to COVID-19**
DIMITRIS BERTSIMAS, *Sloan School of Management at MIT*

15:45 – 16:15 **Ceremony of the Willem R. van Zwet Award & the Jan Hemelrijk Award**
Prize winners will be presented by the juries, followed by a short presentation by the laureates

16:15 – 16:30 Break

16:30 – 17:00 **Curbing the spread: observations, interventions, models, and predictions**
JACCO WALLINGA, *Leiden University* | *RIVM*

17:00 – 17:15 **Final discussion with speakers & Closure**

13:45 – 14:15

CONTAINMENT STRATEGIES FOR COVID-19 USING MOBILITY DATA

Prof. dr. N.V. (Nelly) Litvak

University of Twente | Eindhoven University of Technology

In their response to the COVID-19 outbreak, governments face the dilemma to balance public health and economy. Mobility plays a central role in this dilemma because the movement of people enables both economic activity and virus spread. We use mobility data in the form of counts of travelers between regions, to extend the often-used SEIR models to include mobility between regions. We quantify the trade-off between mobility and infection spread in terms of a single parameter, to be chosen by policy makers, and propose strategies for restricting mobility so that the restrictions are minimal while the infection spread is effectively limited. We consider restrictions where the country is divided into regions, and study scenarios where mobility is allowed within these regions, and disallowed between them. We propose heuristic methods to approximate optimal choices for these regions. We evaluate the obtained restrictions based on our trade-off. The results show that our methods are especially effective when the infections are highly concentrated, e.g., around a few municipalities, as resulting from superspreading events that play an important role in the spread of COVID-19. We demonstrate our method in the example of the Netherlands. The results apply more broadly when mobility data is available.

Joint work with: Martijn Gösgens, Teun Hendriks, Marko Boon, Stijn Keuning, Wim Steenbakkens, Hans Heesterbeek and Remco van der Hofstad.

NELLY LITVAK is Professor of Algorithms for Complex Networks at the University of Twente and Eindhoven University of Technology in the Netherlands. She obtained her PhD in Stochastic Operations research from Eindhoven University of Technology in 2002. Her research interests include random graphs, randomized algorithms, and applications in large networks such as on-line social networks and the World Wide Web.

14:15 – 14:45

INCORPORATING TIME-USE AND HEALTH DATA INTO A DYNAMIC MICROSIMULATION EPIDEMIOLOGICAL MODEL FOR COVID-19

Prof. dr. G. (Gavin) Shaddick

University of Exeter

The need to inform policies and mitigation measures aimed at reducing the spread of the coronavirus highlights the need to understand the complex links between our daily activities and opportunities for the virus to spread. The national lockdowns, and more localised measures, have aimed to reduce the number of contacts between susceptible members of the population and those with the disease. Here, we develop a micro-simulation modelling framework and methods for its computational implementation that brings together epidemiological modelling, urban analytics, spatial analysis and data integration to provide the ability to assess the effects of past interventions and forecast the effects of future policy decisions. This information will be crucial in gaining a greater understanding of the effects of future policy decisions in different areas and within different populations. We demonstrate the use of the model in a case study based on the county of Devon where we compare the effects of different lockdown strategies and present a computationally efficient approach to running complex simulation models of this type.

GAVIN SHADDICK is Professor of Data Science and Statistics and Head of the Department of Mathematics at the University of Exeter. His focus in research is on the theory and application of Bayesian hierarchical models and spatio-temporal modelling in several fields including epidemiology, environmental modelling, and disease progression in rheumatology. Important applications are estimating the hazards of global air quality and research in the power industry. For this, computational techniques that allow the implementation of complex statistical and spatial models are used, with particular attention to the propagation of uncertainty throughout the modelling process.

15:15 – 15:45

FROM PREDICTIONS TO PRESCRIPTIONS: A DATA-DRIVEN RESPONSE TO COVID-19

Prof. dr. D.J. (Dimitris) Bertsimas

Sloan School of Management, MIT

The COVID-19 pandemic has created unprecedented challenges worldwide. Strained healthcare providers make difficult decisions on patient triage, treatment and care management on a daily basis. Policy makers have imposed social distancing measures to slow the disease, at a steep economic price. We design analytical tools to support these decisions and combat the pandemic. Specifically, we propose a comprehensive data-driven approach to understand the clinical characteristics of COVID-19, predict its mortality, forecast its evolution, and ultimately alleviate its impact. By leveraging cohort-level clinical data, patient-level hospital data, and census-level epidemiological data, we develop an integrated four-step approach, combining descriptive, predictive and prescriptive analytics. First, we aggregate hundreds of clinical studies into the most comprehensive database on COVID-19 to paint a new macroscopic picture of the disease. Second, we build personalized calculators to predict the risk of infection and mortality as a function of demographics, symptoms, comorbidities, and lab values. Third, we develop a novel epidemiological model to project the pandemic's spread and inform social distancing policies. Fourth, we propose an optimization model to re-allocate ventilators and alleviate shortages. Our results have been used at the clinical level by several hospitals to triage patients, guide care management, plan ICU capacity, and re-distribute ventilators. At the policy level, they are currently supporting safe back-to-work policies at a major institution and equitable vaccine distribution planning at a major pharmaceutical company, and have been integrated into the US Center for Disease Control's pandemic forecast.

DIMITRIS BERTSIMAS is currently the Boeing Professor of Operations Research, the Associate Dean of Business Analytics at the Sloan School of Management, MIT. He received his MSc and PhD in Applied Mathematics and Operations Research in 1987 and 1988 respectively. His research interests include optimization, machine learning and applied probability and their applications in health care, finance, operations management and transportation. He is the editor in chief of *INFORMS Journal of Optimization* and former editor of *Optimization for Management Science* and *Financial Engineering in Operations Research*.

16:30 – 17:00

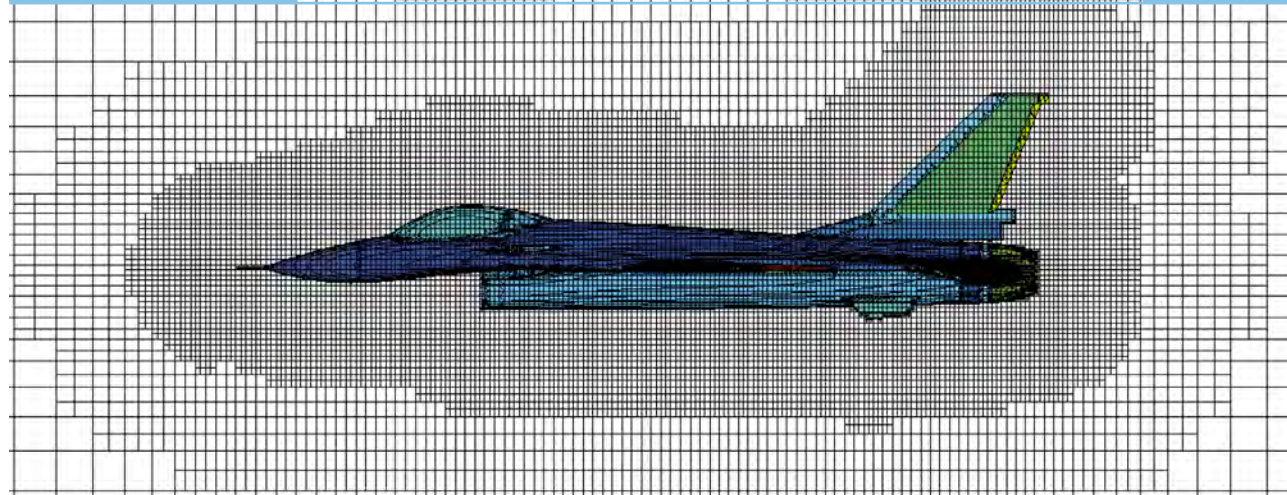
CURBING THE SPREAD: OBSERVATIONS, INTERVENTIONS, MODELS, AND PREDICTIONS

Prof. dr. J. (Jacco) Wallinga

Leiden University | RIVM

The advent of the SARS-CoV-2 virus that causes COVID-19 has elicited an unprecedented research effort to understand the spread, to develop effective interventions, and to predict the possible impact of control measures on the course of the epidemic. Key observations include the number of reported cases over time, and essential interventions include physical distancing and vaccination. The models that are used to describe such observations and analyse the impact of such interventions typically invoke a basic representation of the infection cycle: persons can only be infected after exposure to others who have been infected earlier. Combining such a basic causal model structure with the available observations already allows for predicting the expected impact of control measures. The question is then how this ability to describe, analyse and predict observations can be used to curb the spread.

JACCO WALLINGA holds the chair in Mathematical Modelling of Infectious Diseases at the department of Biomedical Data Sciences of the Leiden University Medical Center (LUMC). As an extraordinary Professor, he works on the estimation and prediction of control measures on the dynamics of infectious diseases. He became head of the RIVM department of Modelling of Infectious Diseases in 2005. He published over 150 articles in international peer-reviewed scientific journals and is a member of the editorial board of *Epidemiology*. He advises the Dutch government and international organizations on vaccination policies.



Bezuinigen op verkiezingen met slimme wiskunde

Slim gebruik van wiskunde kan tot aanzienlijke besparingen leiden. De vele toepassingen van OR zijn hier een voorbeeld van.

Soms zitten die besparingen op plekken die wat minder bekend zijn. Zo kent de numerieke wiskunde de methode van het Adaptive Grid. Bij het berekenen van bijvoorbeeld de luchtstromingen rond een vliegtuig wordt een raster van gelijke vierkanten over het toestel gelegd en wordt voor iedere cel de in- en uitstroom berekend. Het verschil tussen die twee resulteert dan in een over- of onderdruk etc. Die berekening wordt voor iedere cel vele malen herhaald, telkens een tijdstapje verder. Steeds wordt gekeken naar de druk in de naastliggende cellen, immers die bepaalt de mate van in- en uitstroom. Dit kost enorm veel rekentijd, zeker als de cellen en de tijdstappen klein zijn. Bij een Adaptive Grid worden de cellen eerst vrij groot gekozen, vervolgens wordt gekeken in welke cellen zich tijdens de eerste tijdstappen grote veranderingen voordoen. Alleen die cellen worden vervolgens verkleind, de rest wordt ongemoeid gelaten. Dat proces wordt vele malen herhaald: de cellen worden kleiner naarmate er meer verandering in plaats vindt, vandaar de naam Adaptive Grid. Het idee hierachter is dat men geen rekentijd hoeft te verspillen aan die delen van het raster waar de veranderingen klein zijn, men zet de computerkracht hoofdzakelijk daar in waar 'het ertoe doet'.

Toen ik weer eens veel te lang en veel te laat naar CNN keek tijdens de dagen direct na de Amerikaanse verkiezingen van 2020 bedacht ik dat men hier eigenlijk prima dat principe van het Adaptive Grid zou kunnen toepassen. De uitslag wordt, door het in mijn ogen vreemde systeem van 'the winner takes all' kiesmannen, in de praktijk

slechts bepaald door de uitkomsten zo'n 8 tot 10 staten. De toewijzing van de kiesmannen aan Democraten of Republieken in de overige staten ligt immers al van tevoren vast. Een Democraat in Texas kan stemmen tot hij een ons weegt, de kiesmannen van zijn staat gaan toch onveranderlijk naar de Republikeinen. Waarom zou men dan in zulke staten nog iedere vier jaar verkiezingen organiseren, dat is verspilde tijd, geld en moeite. Als er in een staat de uitslag een verschil van meer dan bijvoorbeeld 5% tussen de beide kampen laat zien wordt die staat bij de volgende verkiezing overgeslagen. De vierjaarlijkse verkiezingen worden dan alleen in die staten gehouden die in de praktijk de einduitslag bepalen, zo kan men tot aanzienlijke besparingen komen. Ook hier weer alleen je energie gebruiken daar waar 'het ertoe doet'.

Denk overigens niet dat dit idee van mij revolutionair is. In 1958 las ik een science fiction verhaal ('Franchise' door Isaac Asimov) over de presidentsverkiezing in 2008, toen nog een verre toekomst. Daar gaat het nog veel verder. Op de dag van de verkiezing wordt slechts één enkele kiezer met veel ceremonieel van huis gehaald en naar het stemlokaal gebracht. Daar brengt hij zijn stem uit en dat is het. Amerika was een Electronic Democracy geworden en de computer Multivac, die alles van alle inwoners wist, had bepaald dat deze ene kiezer het perfecte gemiddelde is van alle kiezers in het land: hij vertegenwoordigt daarom in zijn eentje alle kiezers. Dat zou pas een grote besparing zijn! Maar of we zoiets zouden moeten willen?

GERRIT STEMERDINK is eindredacteur van STAtoR.
E-mail: gjsterdink@hotmail.com



ACCUMULATION BIAS

How to handle it ALL-IN

JUDITH TER SCHURE

An estimated 85% of global health research investment is wasted (Chalmers and Glasziou, 2009); a total of one hundred billion US dollars in the year 2009 when it was estimated. The movement to reduce this research waste recommends that previous study results be taken into account when prioritizing, designing and interpreting new research (Chalmers et al., 2014; Lund et al., 2016). Yet any recommendation to increase efficiency this way requires that researchers evaluate whether the studies already available are sufficient to complete the research effort; whether a new study is necessary or wasteful. These decisions are essentially stopping rules – or rather noisy accumulation processes, when no rules are enforced – and unaccounted for in standard meta-analysis. Hence reducing waste invalidates the assumptions underlying

many typical statistical procedures.

Ter Schure and Grünwald (2019) detail all the possible ways in which the size of a study series up for meta-analysis, or the timing of the meta-analysis, might be driven by the results within those studies. Any such dependency introduces *accumulation bias*. Unfortunately, it is often impossible to fully characterize the processes at play in retrospective meta-analysis. The bias cannot be accounted for. Here, we discuss an example accumulation bias process, that can be one of many influencing a single meta-analysis, and use it to illustrate the following key points:

- Standard meta-analysis does not take into account that researchers decide on new studies based on other study results already available. These decisions introduce

accumulation bias because the analysis assumes that the size of the study series is unrelated to the studies within; it essentially conditions on the number of studies available.

- Accumulation bias does not result from questionable research practices, such as publication bias from file-drawering a selection of results. The decision to replicate only some studies instead of all of them biases the sampling distribution of study series, but can be a very efficient approach to set priorities in research and reduce research waste.
- ALL-IN meta-analysis stands for *Anytime, Live and Leading INterim* meta-analysis. It can handle accumulation bias because it does not require a set number of studies, but performs analysis on a growing series – starting from a single study and accumulating as many studies as needed.
- ALL-IN meta-analysis also allows for continuous monitoring of the evidence as new studies arrive, even as new interim results arrive. Any decision to start, stop or expand studies is possible, while keeping valid inference and type-I error control intact. Such decisions can be strategic: increasing the value of new studies, and reducing research waste.

Our example: extreme Gold Rush accumulation bias

We imagine a world in which a series of studies is meta-analyzed as soon as three studies become available. Many topics deserve a first initial study, but the research field is very selective with its replications. Nevertheless, for significant results in the right direction, a replication is warranted. We call this the *Gold Rush* scenario, because after each finding of a positive significant result – the gold in science – some research group rushes into a replication, but as soon as a study disappoints, the research effort is terminated and no-one bothers to ever try again. This scenario was first proposed by Ellis and Stewart (2009) and formulated in detail and under this name by Ter

Schure and Grünwald (2019). Here we consider the most extreme version of the *Gold Rush* where finding a significant positive result not only makes a replication more probable, but even inevitable: the dependency of occurring replications on their predecessor's result is deterministic.

Biased Gold Rush sampling

We assume that we always stop performing studies when we've reached three, so we could possibly perform a meta-analysis that includes $t = 1, 2$ or 3 studies: t indicates the number of studies as well as the timing of the meta-analysis. We summarize the results of individual studies in a single per-study Z-score (z_1 for the first study, z_2 for the second, etc). A significant positive study is shown as z_i^+ ($z_i > z_{\alpha}$ with $z_{\alpha} = 1.96$ for $\alpha = 2.5\%$) and a nonsignificant or negative one as z_i^- . Our *Gold Rush* world consists of the possible study series in the box below.

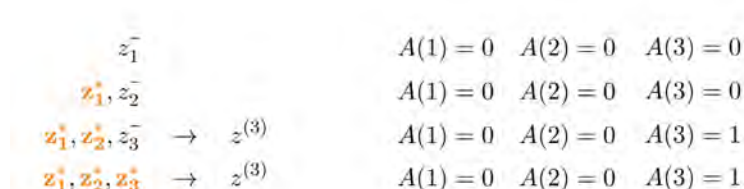
Here $A(t)$ denotes whether we accumulate **and** analyze the t studies: It can be that $A(2) = 0$ and $A(3) = 0$ because we are stuck at one study, but also $A(1) = 0$ because we don't 'meta-analyze' that single study. $z^{(3)}$ indicates the Z-score of a fixed or common effect meta-analysis. The effects of accumulation bias are not limited to fixed-effects meta-analysis (see for example Kulinskaya et al. (2016)), but we use fixed-effects meta-analysis as a simple illustration.

We observe in our *Gold Rush* world below that the study series that are eventually meta-analyzed into a Z-score $z^{(3)}$ are a very biased subset of all possible study series. So we expect these $z^{(3)}$ scores to be biased as well.

The conditional sampling distribution under extreme Gold Rush accumulation bias

Assume that we are in the scenario that only true null effects are studied in our *Gold Rush* world, such that any new study builds on a false-positive result. How large

Gold Rush world



```
numSim.study <- 64000000

Z1 <- rnorm(numSim.study)
Z2 <- rnorm(numSim.study)
Z3 <- rnorm(numSim.study)

# selection based on Gold Rush accumulation bias A(3) = 1
A3 <- which((Z1 > 1.96) & (Z2 > 1.96))
numSim.3series <- length(A3)

calcZmeta <- function(Zs) {
  t <- length(Zs)
  1/sqrt(t)*sum(Zs)
}

# meta Zscores for a random sample of 3-study series
Zmeta3 <- sapply(sample(1:numSim.study, size = numSim.3series), function(i) calcZmeta(c(Z1[i], Z2[i], Z3[i])))

# meta Zscores for a biased sample of 3-study series, biased by GoldRush A(3) = 1
Zmeta3.A3 <- sapply(A3, function(i) calcZmeta(c(Z1[i], Z2[i], Z3[i])))
```

R Code 1 to simulate Figure 1

would the bias be if the three-study series are simply analyzed by standard meta-analysis? We illustrate this by simulating this *Gold Rush* world using R Code 1. A fixed-effect meta-analysis Z-score weights the study estimates by their precision (inverse variance). Here we assume equal variance and large enough (and equal) sample size to estimate the variance without error, in which case the fixed-effects $z^{(3)}$ -score reduces to $1/\sqrt{t} \cdot \sum(zs)$.

Theoretical sampling process

A fixed-effects meta-analysis assumes that if three studies z_1, z_2, z_3 are each sampled under the null hypothesis, each has a standard normal with mean zero and the standard normal sampling distribution also applies for the combined $z^{(3)}$ score. The R code in R Code 1 illustrates this sampling process: First, a large population is simulated of possible first ($Z1$), second ($Z2$) and third ($Z3$) studies from a standard normal distribution. Then in $Zmeta3$ each index i represents a possible study series, such that $c(Z1[i], Z2[i], Z3[i])$ samples an unbiased study

series and $calcZmeta$ calculates its fixed-effects meta-analysis Z-score $z^{(3)}$. So the large number of Z-scores in $Zmeta3$ capture the unbiased sampling distribution that is assumed for fixed-effects meta-analysis $z^{(3)}$ -scores.

Gold Rush sampling process

In contrast, the code resulting in $A3$ selects only those study series for which $A(3) = 1$ under extreme *Gold Rush* accumulation bias. So the large number of Z-scores in $Zmeta3.A3$ capture a biased sampling distribution for the fixed effects meta-analysis $z^{(3)}$ -scores.

Figure 1 plots two histograms of $z^{(3)}$ samples, one with and one without the *Gold Rush* $A(t)$ accumulation bias process, based on $Zmeta3.A3$ and $Zmeta3$ respectively.

We observe in Figure 1 that the theoretical sampling process, resulting in the pink histogram, gives a distribution for the three-study meta-analysis $z^{(3)}$ -scores that is centered around zero. Under the *Gold Rush* sampling process, however, our three-study $z^{(3)}$ -scores do not behave like this theoretical distribution at all. The

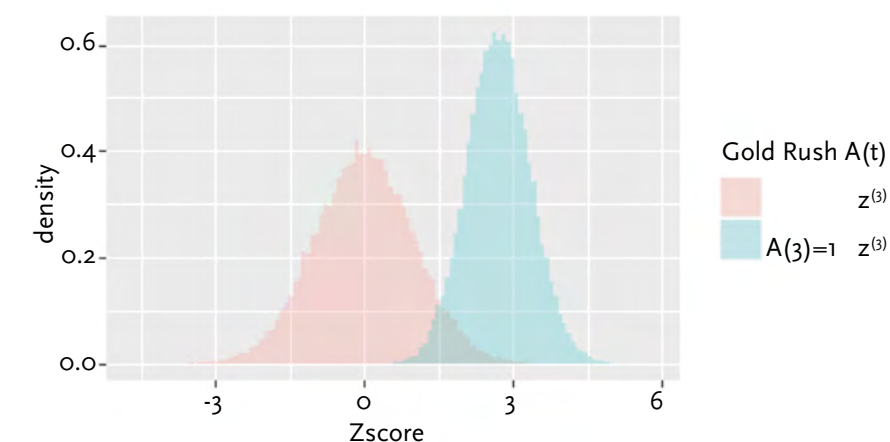


Figure 1. Sampling distributions under the null hypothesis of fixed-effects meta-analysis Z-scores $Z^{(3)}$ of three studies with and without extreme Gold Rush accumulation bias $A(t)$, under the assumption of equally large study sample size and equal variance

blue histogram has a smaller variance and is shifted to the right – representing the bias.

We conclude that we should not use conventional meta-analysis techniques to analyze our study series under *Gold Rush* accumulation bias. Conventional fixed-effects meta-analysis assumes that any three-study summary statistic $Z^{(3)}$ is sampled from the pink distribution in Figure 1 under the null hypothesis, such that the meta-analysis is significant for $Z^{(3)}$ -scores larger than $z_{\alpha} = 1.96$ for a right-sided test with type-I error control $\alpha = 2.5\%$. Yet the actual blue sampling distribution under this accumulation bias process shows that a much larger fraction of series that accumulate three studies will have $Z^{(3)}$ -scores larger than 1.96 than is assumed by the theory of random sampling.

Accumulation bias can be efficient

The steps in the code from R Code 1 that arrive at the sampling distributions in Figure 1 illustrate that accumulation bias is in fact a selection bias. Nevertheless, accumulation bias does not result from questionable research practices, such as publication bias from file-drawing a selection of results. The selection to replicate only some studies instead of all of them biases the sampling distribution of study series, but can be a very efficient approach to set priorities in research and reduce research waste.

By inspecting our *Gold Rush* world a bit closer, we observe that a fixed-effects meta-analysis of three studies

actually *conditions* on this number of studies ($(A(t))$ needs to be $A(3)$ to be 1), and that this conditional nature is what is driving the accumulation bias. In the next section we take the unconditional view.

ALL-IN meta-analysis

Figure 2 shows an example of an ALL-IN meta-analysis. Each of the red/orange/yellow lines represents a study out of the ten separate studies in as many different countries. The blue line indicates the meta-analysis synthesis of the evidence; a live account of the evidence so far in the underlying studies. In fact, *ALL-IN* meta-analysis stands for *Anytime, Live* and *Leading Interim* meta-analysis, in which the *Anytime Live* property assures valid inference under continuously monitoring and the *Leading* property allows the meta-analysis results to inform whether individual studies should be stopped or expanded. This is important to note that such data-driven decisions would invalidate conventional meta-analysis by introducing accumulation bias.

To interpret Figure 2, we observe that initially only the Australian (AU) study contributes to the meta-analysis and the blue line completely overlaps with the red one. Very quickly, the Dutch (NL) study also starts contributing and the blue meta-analysis line captures a synthesis of the evidence in two studies. Later on, also the study in the US, France (FR) and Uruguay (UY) start contributing and the meta-analysis becomes a three-study, four-study and five-study meta-analysis. How many studies contribute to the

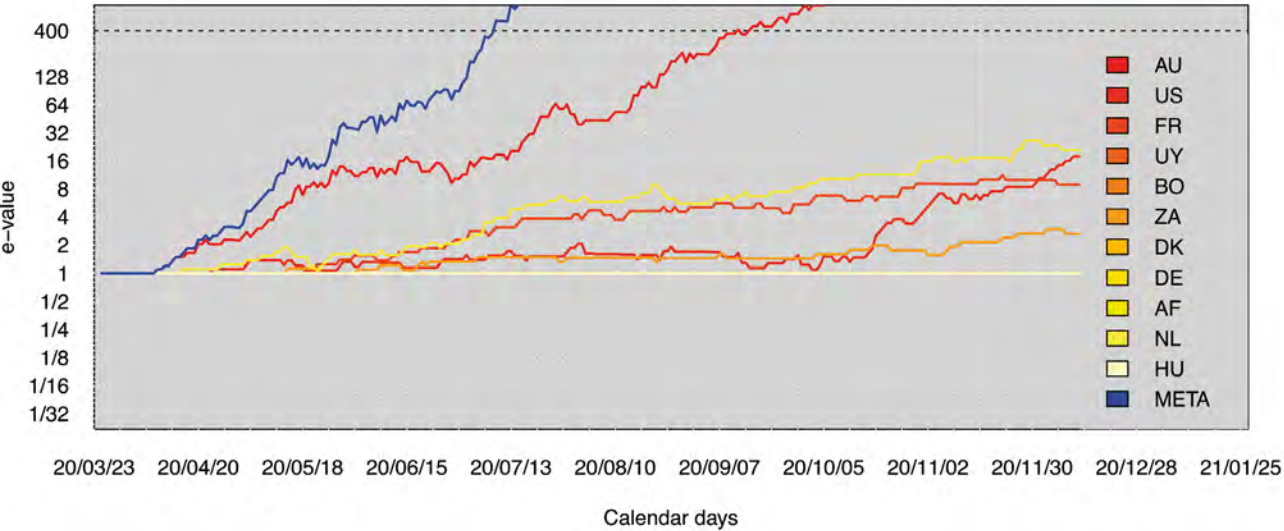


Figure 2. Dashboard of an ALL-IN meta-analysis of between one and eleven studies (with fake data), some of which have not even started recruiting participants in the current status of this dashboard. Note that the y-axis is logarithmic

analysis, however, does not matter for its evidential value. Some studies (like the Australian one) are much larger than others, such that under a lucky scenario this study could reach the evidential threshold even before other studies start observing data. This threshold (indicated at 400) controls type-I errors at a rate of $\alpha = 1/400 = 0.0025$ (details in the final section). So in repeated sampling under the null, the combined studies will only have a probability to cross this threshold that is smaller than 0.25%. In this repeated sampling the size of the study series is essentially random: we can be lucky and observe very convincing data in the early studies, making more studies superfluous, or we can be unlucky and in need of more studies. The threshold can be reached with a single study, with a two-study meta-analysis, with a three-study ... etc, and the repeated sampling properties, like type-I error control, hold on average over all those sampling scenarios (so unconditional on the series size).

ALL-IN meta-analysis allows for meta-analyses with Type-I error control, while completely avoiding the effects of accumulation bias and multiple testing. This is possible for two reasons: (1) we do not just perform meta-analyses on study series that have reached a certain size, but continuously monitor study series irrespective of the current number of studies in the series; (2) we use e-values, that are based on likelihood ratios (Grünwald et al., 2019) instead of raw Z-scores and p -values; we say more on likelihood ratios further below.

Properties averaged over time

Table 1 is inspired by Senn (2014) (different question, similar answer) and represents our extreme *Gold Rush* world of study series. The three study series are very biased, with two or even three out of three studies showing a positive significant effect. But the P_0 column shows that the probability of being in this scenario is very small under the null hypothesis. In fact, most analysis will be of the one-study kind, that hardly have any bias,

and are even slightly to the left of the theoretic standard null distribution. Exactly this phenomenon balances the biased samples of series of larger size.

The bottom row of Table 1 gives the expected values for the number of significant studies per series in the $* \cdot P_0$ column, and the expected value for the total number of studies per series in the $t \cdot P_0$ column. If we use these expressions to obtain the proportion of expected number of significant to expected total number of studies, we get the following:

$$\frac{E_0[*]}{E_0[t]} = \frac{\alpha + \alpha^2 + \alpha^3}{1 + \alpha + \alpha^2} = \frac{\alpha(1 + \alpha + \alpha^2)}{1 + \alpha + \alpha^2} = \alpha$$

The proportion of expected significant effects to expected series size is still α in Table 1 under extreme *Gold Rush* accumulation bias, as it would also be without accumulation bias.

This result is driven by the fact that there is a martingale process underlying this table. A martingale for a series of studies can be thought of as a sequence of statistics, updated after each study, for which the following holds: if the statistic has a certain value after t studies, the conditional expected value of the statistic when the next study is added, so conditional on what is known so far, is equal to the statistic after t studies. The accumulation bias does not affect such statistics when averaged over time (Doob's Optional Stopping Theorem for martingales). You can get a sense of this theorem for *Gold Rush* accumulation bias by deleting the last row for z_1^*, z_2^*, z_3^* from our table and adding two rows for $t = 4$ in its place with z_1^*, z_2^*, z_3^* and either a fourth significant or a nonsignificant study. If you calculate the expected significant effects to expected series size, you will again arrive at α .

Martingale properties drive many approaches to sequential analysis, including the Sequential Probability Ratio Test (SPRT), group-sequential analysis and alpha spending. When applied to meta-analysis, any such inferences essentially average over series size, just like ALL-IN meta-analysis.

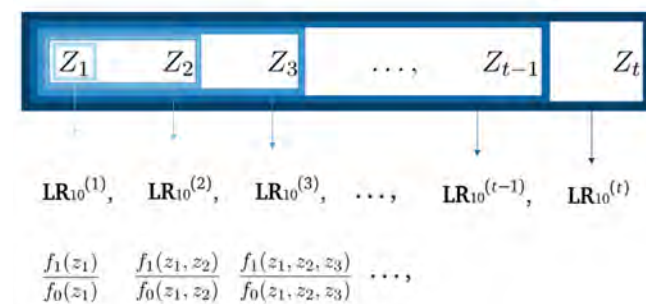
t	*	P_0	$* \cdot P_0$	$t \cdot P_0$
1	z_1^-	0	$1 - \alpha$	0
2	z_1^*, z_2^-	1	$\alpha(1 - \alpha)$	$\alpha(1 - \alpha)$
3	z_1^*, z_2^*, z_3^-	2	$\alpha^2(1 - \alpha)$	$2\alpha^2(1 - \alpha)$
3	z_1^*, z_2^*, z_3^*	3	α^3	$3\alpha^3$
Σ		1	$\alpha + \alpha^2 + \alpha^3$	$1 + \alpha + \alpha^2$

Table 1. Possible study series under extreme *Gold Rush* accumulation bias, with their respective number of significant studies (*) and probabilities (P_0) to occur under the null hypothesis

Multiple testing over time

Just having the expectation of some statistics not affected by stopping rules is not enough to monitor data continuously, as in ALL-IN meta-analysis. We need to account for the multiple testing as well. In that respect, the approaches to sequential analysis differ by either restricting inference to a strict stopping rule (SPRT), or setting a maximum sample size (group-sequential analysis and alpha spending).

ALL-IN meta-analysis takes an approach that is different from its predecessors and is part of an upcoming field of sequential analysis for continuous monitoring with an unlimited horizon. These approaches are called *Safe* for optional stopping and/or continuation (Grünwald et al., 2019) *any-time valid* (Ramdas et al., 2020). Their methods rely on martingales that are nonnegative (Ramdas et al., 2020), specifically the likelihood ratio. For a meta-analysis Z-score, a martingale process of likelihood ratios could look as follows:



The subscript 10 indicates that the denominator of the likelihood ratio is the likelihood of the Z-scores under

the null hypothesis of mean zero, and in the numerator is some alternative mean normal likelihood. The likelihood ratio becomes smaller when the data are more likely under the null hypothesis, but the likelihood ratio can never become smaller than 0 (hence the ‘nonnegative’ martingale). This is crucial, because a nonnegative martingale allows us to use Ville’s inequality (Ville, 1939), also called the universal bound by Royall (1997). For likelihood ratios, this means that we can set a threshold that guarantees type-I error control under any accumulation bias process and at any time, as follows:

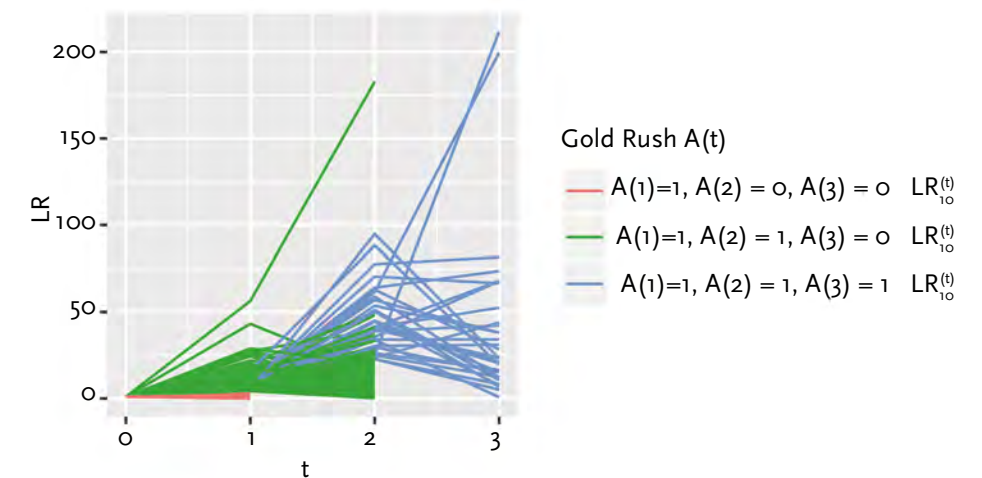
$$P_0 \left[LR_{10}^{(t)} \geq \frac{1}{\alpha} \text{ for some } t = 1, 2, \dots \right] \leq \alpha.$$

The ALL-IN meta-analysis in Figure 2 in fact is based on likelihood ratios like this, and controls the type-I error by the threshold 400 at level $1/400 = 0.25\%$.

The R-Code 2 illustrates that likelihood ratios can also control type-I error rates under continuous monitoring when extreme *Gold Rush* accumulation bias is at play. Just as in our previous simulation, we again assume a *Gold Rush* world with only true null studies and very biased two-study and three-study series. The code in R Code 2 calculates likelihood ratios for the growing study series under accumulation bias. Figure 3 illustrates that still very few likelihood ratio processes ever grow very large.

If we set our type-I error rate α to 5%, and compare our likelihood ratios to $1/\alpha = 20$ we observe that less than $1/20 = 5\%$ of the study series *ever* achieves a value of LR_{10} larger than 20 (R Code 3). The simulated

Figure 3. Samples under the null hypothesis of $LR_{10}^{(t)}$ of one, two or three studies under extreme Gold Rush accumulation bias, under the assumption of equally large study sample size and equal variance.



type-I error is even much smaller than 5% since in our *Gold Rush* world series stop growing at three studies, yet this procedure controls type-I error also in the case none of these series stops growing at three studies, but all continue to grow forever.

The type-I error control is thus conservative, and we pay a small price in terms of power. That price is quite manageable, however, and can be tuned by setting the mean value of the alternative likelihood (arbitrarily set to mean = 1 in the code for calcLR of R Code 2). More on that in Grünwald et al. (2019).

It is this small conservatism in controlling type-I error that allows for full flexibility: There isn’t a single accumulation bias process that could invalidate the inference. Any data-driven decision is allowed. And data-driven decisions can increase the value of new studies and reduce research waste.

Postscript

ALL-IN meta-analysis has been applied during the corona pandemic to analyze an accumulating series of studies while they were still ongoing. Each study investigated the ability of the BCG vaccine to prevent COVID-19, but

data on COVID cases came in only slowly (fortunately). Meta-analyzing interim results and data-driven decisions improved the possibility of finding efficacy earlier in the pandemic.

STAtOR 2020-4 contained an article ‘Nieuwe statistiek voegt wereldwijd corona-onderzoek samen’ that explained the ALL-IN approach in terms of e-values and gambling. This interpretation is closely related to the notion of martingales as a *fair game*. Likelihood ratios are e-values and e-values are (super)martingales and can therefore be interpreted as the betting profit of a fair game.

REFERENCES

The list of references, technical details (e.g. the specific martingale underlying the table) and a more elaborate version of this article are available on the VVSOR-website. The digital article also links to the full R code that runs the simulations and produces the plots.

JUDITH TER SCHURE’s research interests lie in foundations of statistics as well as in applied work. She divides her time between her PhD research on ALL-IN meta-analysis at CWI, freelance statistical consultancy (significantheelp.nl) and board membership (treasurer) of VVSOR. E-mail: mail@judithterschure.nl

```
numSim.study <- 64000 # we're not plotting histograms, so a smaller simulation will do

Z1 <- rnorm(numSim.study)
Z2 <- rnorm(numSim.study)
Z3 <- rnorm(numSim.study)

A1notA2 <- which(Z1 <= 1.96)
A2notA3 <- which((Z1 > 1.96) & (Z2 <= 1.96))
A3 <- which((Z1 > 1.96) & (Z2 > 1.96))

calcLR <- function(Zs) {
  prod(dnorm(Zs, mean = 1)/dnorm(Zs, mean = 0))
}

LR1.A1notA2 <- sapply(A1notA2, function(i) calcLR(Z1[i]))
LR1.A2notA3 <- sapply(A2notA3, function(i) calcLR(Z1[i]))
LR1.A3 <- sapply(A3, function(i) calcLR(Z1[i]))
LR2.A2notA3 <- sapply(A2notA3, function(i) calcLR(c(Z1[i], Z2[i])))
LR2.A3 <- sapply(A3, function(i) calcLR(c(Z1[i], Z2[i])))
LR3.A3 <- sapply(A3, function(i) calcLR(c(Z1[i], Z2[i], Z3[i])))
```

R Code 2 to create Figure 3

```
> typeIerrorLR <- mean(c(LR1.A1notA2,
+                        pmax(LR1.A2notA3, LR2.A2notA3),
+                        pmax(LR1.A3, LR2.A3, LR3.A3))
+                        > 20)
> typeIerrorLR
[1] 0.001390625
```

R Code 3 to calculate type-I error probability for continuous testing $LR_{10}^{(t)}$ averaged over series size t under extreme Gold Rush accumulation bias



Generaal Eisenhower spreekt tot de luchtlandingstroepen aan de vooravond van D-Day, 5 juni 1944.
Foto: U.S. Army Signal Corps

Radicaal

De coronacrisis en de daaruit volgende lockdown zijn absoluut beroerd, maar er zitten ook een paar onverwachte voordelen aan. Een ervan is dat ik meer boeken kan lezen, waaronder twee die net voor de corona-uitbraak verschenen: *Uncharted*¹ (Heffernan) en *Radical Uncertainty*² (Kay & King). Beide boeken behandelen het omgaan met onzekerheid in besluitvorming. Beide boeken brengen ook een vergelijkbare boodschap: de toegevoegde waarde van modellen is in (radicaal) onzekere tijden miniem. Die conclusie lijkt op het eerste gezicht logisch, maar onderschat wat mij betreft de waarde van modellen.

Onzekerheid

Om te beginnen: wat is radicale onzekerheid nu precies? Onzekerheid kan op heel verschillende manieren worden gecategoriseerd, maar een driedeling die ik zelf veel gebruik is de volgende:

1. Onzekerheid waarvan de kansverdeling exact bekend

is. Denk aan het opgooien van een munt of dobbelsteen, een roulettespel, et cetera.

2. Onzekerheid waarvan een empirische kansverdeling kan worden geschat. Daar valt bijvoorbeeld het schatten van de aankomstverdeling van patiënten bij de Spoedeisende Hulp onder.
3. Onzekerheid die alleen in subjectieve kansen kan worden uitgedrukt, bijvoorbeeld de kans op succes van de introductie van een nieuw product als een iPad, of een virusuitbraak.

De laatste categorie van onzekerheid wordt door Kay & King als radicaal betiteld, een staat van onzekerheid die wordt gekenmerkt door niet-stationariteit en reflexiviteit. In dergelijke situaties moet volgens de auteurs het gebruik van modellen gebaseerd op exacte of empirische kansverdelingen worden afgeraden. Hierin volgen ze Nassim Taleb die in *The Black Swan*³ al waarschuwt tegen het toepassen van naïeve en vereenvoudigde statistische modellen in complexe domeinen als de economie en aandelenmarkten. Kay & King geven veel voorbeelden

om hun stelling te onderbouwen, onder meer het falen van macro-economische modellen omdat recessies, oorlogen en technologische doorbraken niet te voorspellen zijn, waardoor ze een totaal verkeerd beeld geven van de toekomstige ontwikkeling van macro-economische grootheden als GDP. In plaats van het gebruik van kansmodellen in deze radicale omstandigheden wijzen Kay & King op de waarde van *narratives* (verhalen). Narratives helpen ons beslissingen te nemen waarin waarschijnlijkheidsoordelen niet haalbaar zijn. Ze helpen bij het beantwoorden van de vraag 'wat is er aan de hand?' en bij het concreet maken van abstracte begrippen. Een narrative hoeft niet noodzakelijk waar te zijn. Een fictieve narrative kan informatie en inzichten overbrengen die ons helpen om door een onzekere situatie te navigeren.

Modellen zijn geen orakels

Nu wil ik niets afdoen aan het nut en de waarde van narratives. Maar het in de ban doen van modellen in radicaal onzekere situaties berust mijns inziens op een misverstand. Dat misverstand komt voort uit de manier waarop modellen vaak worden ingezet, en waarop niet alleen Kay & King, maar ook anderen terecht kritiek hebben.

Te vaak worden modellen gezien als orakels. Mensen verwachten van een model dat het hét antwoord geeft, dat het een beslissing voorstelt. Maar dat is de functie van een model helemaal niet. Modellen zijn abstracties van de werkelijkheid. Het doel van een model is inzicht te verschaffen, zodat je een situatie beter leert begrijpen en een weloverwogen beslissing kunt nemen. Het ontwikkelen van modellen is geen doel op zich. Het is een ontdekkingsreis waarin modellen steeds verder evolueren en uiteindelijk leiden tot inzichten die je er niet zelf in hebt gestopt. Een mooi voorbeeld is de manier waarop mijn team en ik jaren geleden TNT Express (nu onderdeel van FedEx) met behulp van modellen ondersteunden bij het begin van de financiële crisis. Door die crisis viel het vrachtvolume enorm terug, en TNT stond voor de vraag of het een groot aantal verbindingen in het Europese weg- en luchtvrachtnetwerk zou sluiten of niet. Dat zou de kosten terugdringen, maar zou mogelijk ook leiden tot derving van inkomsten en verlies van klanten. Door de inzet van modellen kon de impact op kosten, inkomsten en service van verschillende alternatieven inzichtelijk worden gemaakt voor diverse vrachtvolume-scenario's.

Het waren niet de uitkomsten van de modellen, maar de verkregen inzichten die de toenmalige board van TNT in staat stelden robuuste beslissingen te nemen waardoor TNT weer in de zwarte cijfers belandde. Dat was een unicum, omdat dit type besluit tot op dat moment doorgaans werd genomen op basis van overtuigingen en narratives.

Een goed model is altijd nuttig

Iedereen kent de uitspraak van George Box: 'All models are wrong, but some are useful'⁴. Het nut van modellen zit in het verkrijgen van inzichten op basis waarvan weloverwogen beslissingen kunnen worden genomen. Een model is echter noch een voldoende, noch een noodzakelijk voorwaarde om een probleem succesvol aan te pakken. Veel hangt af van de omstandigheid (hoe radicaal ook) en de modelvaardigheid (het ontwerpen van een model is een vak!). Een goed toegepast model kan ook in radicaal onzekere tijden bijzonder nuttig zijn. Het uitsluiten van modellen lijkt mij dan ook een vergissing. Het is niet een kwestie van narratives óf modellen, maar van narratives én modellen. Narratives kunnen namelijk worden omgezet in modellen of scenario's waarmee kan worden gerekend. Daarmee kan de geloofwaardigheid van narratives worden gekwantificeerd, en daaruit vloeien inzichten voort die op hun beurt weer tot nieuwe narratives en modellen leiden. Om Dwight D. Eisenhower⁵ te parafraseren: 'models are useless, modelling is indispensable'.

1. Heffernan, M. (2020). *Uncharted; How to map the future*. Export/Airside.
2. Kay, J., & King, M. (2020). *Radical Uncertainty; Decision-making beyond the numbers*. W. W. Norton Company.
3. Taleb, N. N. (2008). *The Black Swan; The Impact of the highly improbable*. Random House.
4. https://en.wikipedia.org/wiki/All_models_are_wrong
5. https://en.wikiquote.org/wiki/Dwight_D._Eisenhower, 'In preparing for battle, I have always found that plans are useless but planning is indispensable.'

De volgende boeken staan op mijn leeslijst:
Shiller, R. J. (2019). *Narrative Economics; How stories go viral and drive major economics events*. Princeton University Press.
Madsbjerg, C. (2019). *Sensemaking; What makes human intelligence essential in the age of the algorithm*. Hachette.

JOHN POPPELAARS, Directeur LITIC B.V.
E-mail: john.poppelaars@litic.com



HOE LOOP JE EEN MARATHON?

Of het nu gaat om korte sprintjes of het lopen van marathonaftanden, uiteindelijk wil je eruit halen wat erin zit en optimaal presteren of genieten. Tijdens een marathon moet een loper zijn krachten goed verdelen over de gehele afstand.

Te snel starten kan zorgen voor voortijdige vermoeidheid, terwijl te traag starten niet tot een toptijd zal leiden. Aan het einde van de rit is het fijn de beschikbare energetische reserves optimaal te hebben gebruikt. Dit proces noemen we ook wel 'pacing' (tempo) strategie.

BRYAN VAN INGEN

Een marathon loop je niet zomaar, hier gaat veel training en voorbereiding aan vooraf. Het zou zonde zijn als een loper daarna langzamer dan gewenst, of zelfs helemaal niet kan finishen. Niet iedereen kan een marathon lopen zoals Eluid Kipchoge of Lornah Kiplagat maar elke sporter wil wel bij elke race zoveel mogelijk uit zichzelf halen. Daarbij is het belangrijk dat elke loper – recreant of professional – met een plan aan de start verschijnt. Tijdens dit onderzoek is gekeken naar welke pacingstrategieën het vaakst door marathonlopers worden ingezet.

Hardlopen wordt steeds populairder. In tijden van lockdowns en quarantaines is hardlopen een van de weinige sporten die men nog ongehinderd kan blijven beoefenen. Ook recreatieve hardlopers pakken de trainingen serieus aan. Veel hardlopers hebben uiteindelijk één ultiem doel: de marathon. Steeds meer sporters leggen

zich toe op de magische afstand van 42,2 kilometer. Eén vraag is voor zowel de recreatieve als de professionele sporter heel belangrijk: 'Hoe moet een loper zijn race indelen om zo snel mogelijk de finish te bereiken?'

Data verzamelen

Le Champion, de organisator van de Amsterdamse Marathon, plaatst de uitslagen van de marathon elk jaar online. Sport Data Valley heeft vervolgens in samenwerking met Amsterdam Data Science de data van 2016 tot en met 2019 samengevoegd voor hun Sport Data Challenge 2019. Deze data zijn voor het onderzoek gebruikt.

De data bestaan grotendeels uit tussentijden van de lopers. Van alle hardlopers die meededen aan de Amsterdamse

marathon tussen 2016 en 2019 is bekend wat hun eindtijd is vergeleken met het startschot, en vergeleken met de daadwerkelijke passage van de startlijn. Hierbij zijn tijdsmetingen van de tussentijden van de lopers na elke vijf kilometer ook bekend. Ook weten we van elk individu het geslacht en de leeftijdscategorie waarin de loper valt. Deze leeftijdscategorieën beginnen bij de senioren (18 tot en met 34 jaar oud) en volgen hierna steeds in intervallen van vijf jaar. In totaal zijn de tijden van 49.245 lopers over de vier jaar gebruikt voor het onderzoek, 37.441 mannen en 11.804 vrouwen.

Voordat we de strategieën van de lopers onderzoeken, kijken we eerst naar de eindtijden van de deelnemers. Hier zien we dat, zoals was te verwachten, mannen met een gemiddelde eindtijd van 3 uur en 56 minuten, gemiddeld bijna een half uur sneller lopen dan vrouwen. Vrouwen hadden gemiddeld 4 uur en 21 minuten nodig om hun 42,2 kilometer af te leggen. Bij de leeftijdsgroepen vinden we wel iets opmerkelijks. Vooraf zou men verwachten dat de gemiddelde snelheid af zou nemen naar mate de leeftijd hoger wordt. Echter, wat bleek, de lopers van 35 tot en met 39 jaar deden niet onder voor de jongere lopers. Ze hadden een gemiddelde snelheid van ongeveer 11 kilometer per uur. In de groepen ouder dan 40 vonden we wel steeds langzamere tijden.

De pacingstrategieën

Voor het onderzoek is de pacing van een marathonloper gedefinieerd als het verloop van de gemiddelde snelheden over de intervallen waarvan de tussentijden bekend zijn. Om alle lopers op eenzelfde manier te kunnen analyseren, wordt er voor het bepalen van de strategieën niet gekeken naar de ware snelheden per interval. In plaats hiervan maken we gebruik van het verschil tussen de gemiddelde snelheid van een loper in ieder interval en hun gemiddelde snelheid over de gehele race. Door deze nieuwe variabele te gebruiken, zou er qua strategie geen verschil moeten zijn tussen iemand die constant 20 km per uur loopt en iemand die constant 8 km per uur loopt.

De pacingstrategieën worden gevonden aan de hand van een clusteranalyse. Hierbij is er een k-means-algoritme gebruikt op de nieuwe variabelen. Na de data genormaliseerd te hebben, kiest dit algoritme meerdere malen k willekeurige gemiddelden (één voor ieder cluster) en bepaalt vervolgens aan welk cluster een persoon toegewezen moet worden zodat de variantie tussen de variabelen binnen de clusters zo laag mogelijk ligt. Zodra iedere loper aan een cluster is toegewezen, worden k nieuwe gemiddelden bepaald op basis van de gemaakte clusters en

wordt iedere loper opnieuw aan een cluster toegewezen. Dit proces wordt 25 keer herhaald om te verzekeren dat de gevonden clusters een zo laag mogelijke intra-cluster variantie hebben, ongeacht welke gemiddelden gekozen werden bij de initiatie van het algoritme.

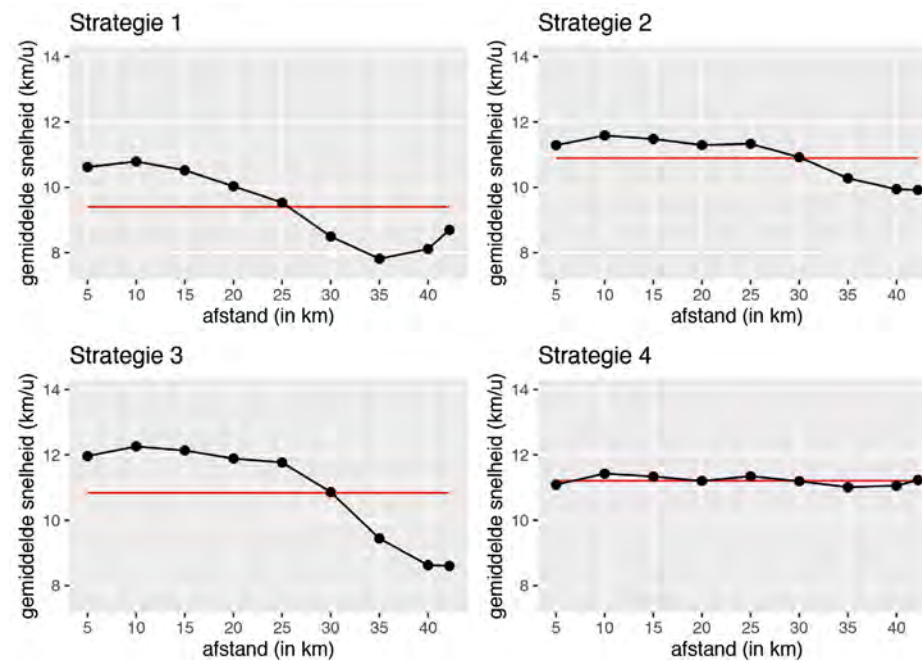
De keuze voor het aantal pacingstrategieën waarmee wij verder werken is afhankelijk van de mate waarin de strategieën onderling van elkaar verschillen. We beginnen met twee strategieën en vervolgens verhogen we dit meermaals met één om de effecten te bekijken. Hierin wordt er gezocht naar grote verschillen tussen de strategieën die gevonden werden. Verschillen zoals lopers in cluster A starten even hard, maar hebben een eindsnelheid van een halve kilometer per uur langzamer dan lopers in cluster B worden niet als onderscheidend genoeg gevonden om mee te nemen. Het blijkt dat als we werken met 4 clusters, en dus 4 strategieën, we het meeste aantal strategieën hebben waar een duidelijk verschil tussen te zien is. De gemiddelde opbouw van de race van lopers in ieder cluster, is weergegeven in figuur 1.

Op basis van wat we in figuur 1 kunnen zien, zijn de vier pacingstrategieën als volgt gedefinieerd:

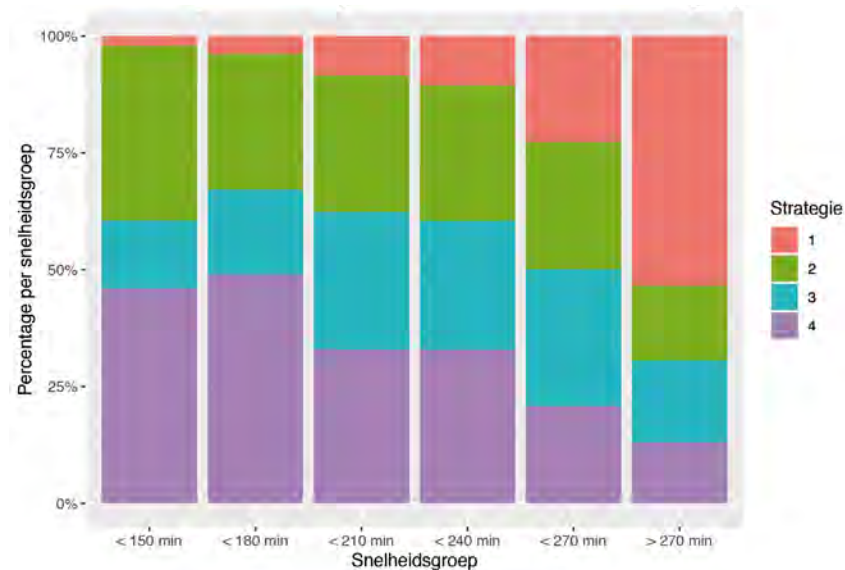
1. De loper start snel, maar gedurende de race verliest hij een groot deel van zijn snelheid. Aan het eind kan de loper er nog wel een korte versnelling uit halen.
2. De loper probeert constant te lopen, de tweede helft zakt het tempo echter wel enigszins in.
3. De loper begint snel, maar in de tweede helft van de race stort de snelheid compleet in.
4. De loper houdt een constante snelheid aan gedurende de race.

In tabel 1 staat weergegeven door hoeveel marathonlopers iedere pacingstrategie gevolgd werd. We zien hier dus dat het grootste deel van de lopers op zijn minst probeert zijn race constant te lopen. Bij de ruim 18.000 lopers die volgens de tweede strategie lopen, zien we dat ze dit helaas niet hebben kunnen volhouden tot de finish.

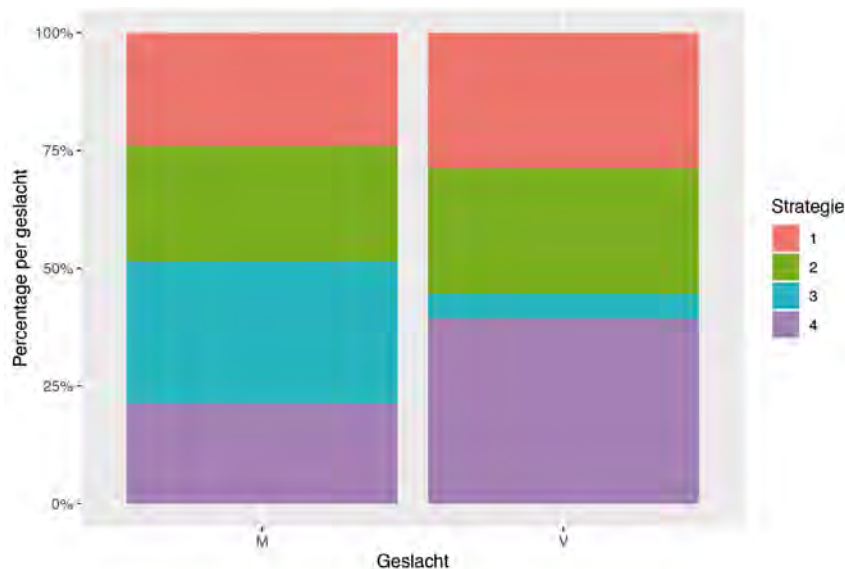
Aangezien we nu onze strategieën hebben, kunnen we kijken hoe de lopers uit verschillende groepen zich verdelen over deze strategieën. Om te kunnen zeggen welke strategie het best lijkt te werken, zullen we de deelnemers eerst in groepen indelen aan de hand van hun eindtijd. Bij deze verdeling beginnen we bij een eindtijd van 2 uur en maken we voor ieder half uur een nieuwe groep. De langzaamste groep zal de groep zijn waarin we iedereen indeelden die langzamer dan viereneenhalf uur deed over de marathon. Natuurlijk zeggen we niet dat dit slechte tijden zijn, maar door de deelnemers op deze manier te verdelen, leggen we de focus bij de snelste groepen en kan er goed gezien worden of daar grote verschillen lig-



Figuur 1. De gemiddelde opbouw van een marathon per pacing-strategie. De rode lijn representeert de gemiddelde snelheid van alle lopers die de strategie volgen



Figuur 2. De verdeling van de pacing-strategieën over de snelheidsgroepen, weergegeven als percentage van elke groep



Figuur 3. De verdeling van de pacing-strategieën per geslacht, weergegeven als percentage van elk geslacht

STRATEGIE	1	2	3	4
AANTAL MARATHONLOPERS	7.580	18.434	5.481	17.750
PROCENTUEEL VAN TOTAAL	15,4%	37,4%	11,1%	36,0%

Tabel 1: Het aantal marathonlopers per pacingstrategie met percentage van totaal.

gen. In figuur 2 zien we per snelheidsgroep de verdeling over de strategieën.

In de figuur is duidelijk te zien dat naarmate de eindtijden sneller worden, de lopers vaker strategie 2 of 4 lijken te volgen, terwijl bij langzamere lopers strategie 1 juist steeds dominanter wordt. Hieruit kan geconcludeerd worden dat snelle lopers vaker constanter lopen. Het is wel belangrijk dat we ons hierbij realiseren dat dit onderzoek niet zegt wat de oorzaak en wat het gevolg is in dit verband. We kunnen niet op basis van deze resultaten zeggen of mensen die een eindsprint proberen te doen, daardoor langzamer lopen, of dat mensen die langzamer lopen vaker nog energie over hebben voor een eindsprint.

In figuren 3 en 4, zien we op dezelfde manier de verdelingen per geslacht en leeftijdscategorie respectievelijk. Hetgeen dat het meest opvalt in figuur 3, is dat bij mannen strategie 3 dominant is, terwijl vrouwen vaker strategie 4 volgen. Mannen mogen dan dus wel sneller lopen, ze schatten zichzelf ook vaak iets te goed in waardoor ze instorten, terwijl vrouwen op een constant tempo naar de finish gaan.

Bij de verschillende leeftijdsgroepen valt op dat strategie 4 steeds minder vaak voorkomt bij hogere leeftijden. Het lijkt er dus op dat oudere deelnemers steeds meer moeite krijgen met het aanhouden van een constante snelheid. Dit gaat gepaard met een groei bij strategie 2. Oudere lopers hebben dus steeds meer moeite met het vasthouden van een constant tempo, maar ze storten niet

compleet in. Het lukt lijkt de ouderen nog steeds goed te lukken om hun snelheid vast te houden, ze verzwakken alleen licht in de tweede helft van de marathon.

Conclusie

Geslacht, leeftijd en de gemiddelde snelheid van een loper hebben alle drie samenhang met de pacingstrategieën van marathonlopers. Over het algemeen kunnen we zien dat lopers met een snellere eindtijd, en dus een hogere gemiddelde snelheid, vaker een constante pacing aanhouden. Bij verschillende leeftijden zien we dat oudere deelnemers vaker moeite hebben het vasthouden van hun snelheid en daardoor gedurende de tweede helft van de race vertragen. Bij de geslachten zien we dat vrouwen vaker constanter lopen dan mannen, die een neiging hebben om te hard te beginnen en vervolgens instorten.

BRONNEN

Dataset: Sport Data Valley, (2019) , Github repository. <https://github.com/sportdatavalley/sport-data-challenge-2019>

BRYAN VAN INGEN studeert Toegepaste Wiskunde aan de Hogeschool van Amsterdam. Met drie studiegenoten analyseerde hij in opdracht van Sport Data Valley, data van de Amsterdamse Marathon. Dit artikel is een samenvatting van het onderzoek. Zie voor een kort filmpje <https://bit.ly/3qC7LNH>. E-mail: Bryan@famingen.nl



Figuur 4. De verdeling van de pacing-strategieën over de leeftijdscategorieën, weergegeven als percentage van elke categorie

HET BALLETTJE KAN RAAR ROLLEN

Foto: Alejandro Garay via Pixabay

Opnieuw veel media-aandacht voor de uitkomst van een lottotrekking die de schijn wekt van frauduleus handelen. In de trekking van de Zuid-Afrikaanse loterij Powerball op dinsdag 1 december 2020 waren de getallen 5, 6, 7, 8, 9 en 10 de getrokken lottogetallen. De reguliere lottogetallen 5, 6, 7, 8 en 9 waren getrokken uit de getallen 1 tot en met 50 en het bonusgetal (Powerball) 10 uit de getallen 1 tot en met 20. Het aantal winnaars van de jackpot was 20 en elk van hen won 5,7 miljoen Zuid-Afrikaanse rand (ongeveer 310 duizend euro), terwijl er 79 deelnemers waren die de reguliere lottogetallen 5, 6, 7, 8 en 9 goed hadden maar niet het Powerball-getal 10 en elk van deze deelnemers won 6283 Zuid-Afrikaanse rand (ongeveer 340 euro). Het aantal winnaars lijkt groot maar bedacht moet worden dat veel deelnemers hun getallen niet random kiezen. Deelnemers gebruiken geboortedata, geluksgetallen, rekenkundige rijen, geometrische patronen, etc. om hun getallen te kiezen. Invullen van het rijtje 1-2-3-4-5-6 is het domste wat je in de lotto kunt doen, niettemin blijkt uit onderzoek dat dit rijtje het meest ingevulde rijtje is in vele lotto's. Voor dit rijtje is de kans dat de zes lottoballetjes erop vallen weliswaar hetzelfde als voor elk ander rijtje, maar in het uitzonderlijk onwaarschijnlijke geval dat de zes lottoballetjes op het rijtje 1-2-3-4-5-6 zouden vallen, moet de jackpot door vele deelnemers worden gedeeld.

Welke kuddedieren deelnemers aan loterijen kunnen zijn in het kiezen van hun lottogetallen wordt ook fraai geïllustreerd door het volgende ware gebeuren. Eén van de meest bekeken dramaserie ooit op de Amerikaanse televisie was de serie *Lost*, die zes seizoenen speelde van

2004-2010. In de episode 'Numbers' van 2 maart 2005 koos de hoofdrolspeler Hurley Reyes de zes getallen 4, 8, 15, 16, 23 en 42 voor een fictieve trekking van de Mega Millions Lotto en won daarmee zogenaamd de jackpot van \$114 miljoen. Binnen de kortste tijd vormden deze zes getallen het meest populaire rijtje van lottogetallen na de rijtjes 1-2-3-4-5-6 en 7-14-21-28-35-42. Dit feit kreeg veel mediapubliciteit na de Mega Millions trekking van 4 januari 2011 met een record jackpot van \$380 miljoen. Wat gebeurde er toen? Uit de getallen 1 tot en met 50 werden de zes getallen 4, 8, 15, 25, 42 en 47 getrokken met 42 als het Mega getal dat ook in de tv serie als Mega getal was gekozen. De drie laagste getallen (4-8-15) en het Mega getal (42) kwamen overeen met Hurley's keuze in de fictieve trekking in de tv serie. In totaal hadden 41.763 deelnemers deze vier getallen goed voorspeld en wonnen daarmee elk \$150.

Direct na de Zuid-Afrikaanse lottotrekking van 1 december 2020 werden de sociale media door veel verbijssterde Zuid-Afrikanen gebruikt om te beweren dat de lottouitslag frauduleus was. '*Lotto exposing themselves that they are a scam*', schreef een twitteraar, en een ander twitterde '*Absolutely no way in hell that's a coincidence*.' Sommigen riepen zelfs op tot een gerechtelijk onderzoek, vergelijkbaar met het onderzoek naar vermeende corruptie tijdens de negenjarige ambtsperiode van ex-president Jacob Zuma. De Nationale Loterijcommissie van Zuid-Afrika, die het spel reguleert, zei dat de combinatie van zes opeenvolgende cijfers ongekend was en beloofde een commissie in te stellen die een evaluatie zal uitvoeren.

De woordvoerder van de loterij voegde daar aan toe 'als er iets is misgegaan, zullen we dat melden'. Overigens wordt de uitkomst van een trekking in de Zuid-Afrikaanse Powerball loterij niet op ouderwetse wijze bepaald door een mechaniek met ronddraaiende balletjes in aanwezigheid van een notaris, maar bepaald met een toevalsgenerator op een computer.

De Zuid-Afrikaanse universiteiten hebben gelukkig goede kansrekenaars die verstand hebben van data-analyse en deze zullen ongetwijfeld met de conclusie komen dat alle ophef rond de trekkingsuitslag een storm in een glas water is. Het 'uitzonderlijke' zal altijd wel ergens een keer optreden als je maar lang genoeg wacht. Het is waar dat de kans op de getallen 5, 6, 7, 8, 9 en 10 in één enkele trekking van de Zuid-Afrikaanse loterij uitzonderlijk klein is, deze kans is 1 op de 42.375.200, maar bedacht moet worden dat er vele loterijen op de wereld zijn en elke loterij gemiddeld twee trekkingen per week heeft. Laten we eens uitgaan van wereldwijd 200 loterijen elk met twee trekkingen per week. Voor het gemak van de berekeningen nemen we aan dat elk van deze loterijen een 6/42 loterij is waarin zes verschillende getallen random worden getrokken uit de getallen 1 tot en met 42. Deze aanname heeft geen invloed op de conclusie die we uit de berekeningen zullen trekken. De 'uitzonderlijke' gebeurtenis waarop we ons zullen richten is de gebeurtenis *E* dat een zestal opeenvolgende getallen worden getrokken, dit mag dus elk van de 37 mogelijke combinaties van zes opeenvolgende getallen uit de getallen 1 tot en met 42 zijn. De kans dat de gebeurtenis *E* zal optreden in één enkele trekking in één enkele 6/42 loterij is 1 op de 141.778. Een eenvoudige berekening geeft dan dat over een periode van 1 jaar ergens op de wereld in een 6/42 loterij zes opeenvolgende getallen worden getrokken gelijk is aan 13,6%, over een periode van 2 jaar is deze kans gelijk aan 25,4% en over een periode van 5 jaar is de kans 52,0%. Deze berekeningen plaatsen de 'onmogelijk' geachte gebeurtenis in de Zuid-Afrikaanse loterij in een ander, meer passend licht. Maar ja, daarmee verwerf je geen media-aandacht, wel met een schreeuwende kop in de trant van 'Is deze loterij uitslag met een kans van 1 op de 42 miljoen doorgestoken kaart?'

HENK TIJMS is emeritus-hoogleraar operations research aan de Vrije Universiteit en auteur van diverse leerboeken over operations research en kansrekening.
E-mail: h.c.tijms@xs4all.nl

'Met analytics de wereld verbeteren'

Joaquim dos Santos Gromicho is benoemd tot hoogleraar Business Analytics aan de Amsterdam Business School en is daarmee de nieuwste aanwinst van de Universiteit van Amsterdam (UvA). Deze van origine Portugese hoogleraar combineert al jaren met veel plezier zijn werk in het bedrijfsleven bij ORTEC met een universitaire loopbaan. Prof. dr. Joaquim dos Santos Gromicho – hoofdredacteur van *STATOR* – vertelt vol overgave over zijn emigratie naar Nederland, zijn passie voor de wetenschap en de potentie van big-data analytics.

Na de afronding van zijn master aan de Universidade de Lisboa werd Joaquim Gromicho gevraagd om twee Erasmusdocenten die gastcolleges gaven op de universiteit een rondleiding te geven door de stad. Al snel bleek dat deze docenten de universiteit met een dubbele agenda bezochten. 'Ze waren op zoek naar talent. In Nederland was het in die tijd voor goede studenten niet aantrekkelijk om te promoveren. Daarom vroegen ze of ik wilde promoveren bij Alexander Rinnooy Kan.'

Twee nieuwe liefdes

Zo gebeurde het dat Gromicho van '91 tot '95 op de Erasmus Universiteit in Rotterdam aan zijn promotieonderzoek werkte. Dat leverde hem meer op dan alleen een PhD. 'In die tijd liep ik ook een knappe, Nederlandse blondine tegen het lijf, die inmiddels mijn vrouw is.' En om de keuze voor Nederland nóg makkelijker te maken, kreeg Gromicho ook een aanbod van ORTEC, 's werelds grootste leverancier van wiskundige optimalisatiesoftware. 'Zij vroegen of ik eens langs wilde komen. Ik voerde leuke gesprekken en kwam erachter dat het bedrijfsleven niet zo hectisch is als ik dacht. Aan het einde van de dag kreeg ik een contract onder mijn neus geschoven.'

Inmiddels is Gromicho naar eigen zeggen al bijna

25 jaar 'getrouwd' met ORTEC. Ooit begon hij daar als software-architect en inmiddels is hij als Scientific and Education Officer verantwoordelijk voor het onderhouden van de relaties met kennisinstellingen. 'Juist doordat wij zulke goede banden onderhouden met het onderwijs in Nederland hebben we het voorrecht dat de beste studenten voor en met ons willen werken. De interne sfeer hier is ook vergelijkbaar met die van een universiteit. We doen in wezen ook bijna hetzelfde: wetenschappelijk onderzoek. Alleen ligt bij ORTEC de nadruk niet op publicaties, maar op commercie. Wij verkopen innovatie.'

Verlangen naar de wetenschap

Toch begon het na vijf jaar bij ORTEC te kriebelen bij Gromicho. 'Ik miste de academische wereld; het werken aan publicaties en het vrijer indelen van mijn tijd en aandacht. Toevallig nam in die tijd een oud-collega in Hongkong contact met mij op. Hij wilde per se mijn hulp bij het afronden van een paper.' Gromicho pakte zijn koffers, nam verlof op en ging voor zes weken naar Hongkong. 'Bij terugkomst vertelde mijn baas, hoogleraar bij de VU, dat je in Nederland prima een academische loopbaan kunt combineren met een carrière in het bedrijfsleven. Dat is in Portugal wel anders. Er ging een wereld voor mij open. Ik ging in deeltijd op de VU aan de slag, waar ik in 2010 werd benoemd tot bijzonder hoogleraar Applied Optimisation in Operations Research.'

Nadat Marc Salomon, voorzitter van de Amsterdam Business School (ABS), Gromicho een aanbod deed dat hij 'niet kon weigeren', maakte hij de overstap naar de UvA. 'Marc belde me, omdat hij mijn hulp wilde bij het vormgeven van een opleiding binnen de Amsterdam Business School. Een fantastische uitdaging. En hij had Dick den Hertog al zover gekregen om ook naar de UvA te komen. Dick was lange tijd hoogleraar aan de Universiteit van Tilburg, een autoriteit binnen mijn vakgebied en een van mijn persoonlijke helden. Toen ik hoorde dat hij zich ook had aangesloten, was ik direct om.' Inmiddels werkt Gromicho 4 dagen bij ORTEC en één dag bij de ABS, als hoogleraar Business Analytics. Zijn onderwijs focust zich op het gebied van *big data-analytics*, met de nadruk op (optimale) besluitvorming.

De kracht van het vak

De passie voor zijn vakgebied is voelbaar bij Gromicho. 'Al van jongs af aan wilde ik de wereld redden. Om dat te bereiken, zullen we ons gedrag drastisch moeten veranderen. Het analyseren van big-data kan daar een grote rol in spelen.' ORTEC en de UvA delen die persoonlijke overtuiging van Gromicho, wat voor hem de zingeving van zijn werkzaamheden vergroot. 'Bij ORTEC geloven we dat wiskunde kan bijdragen aan duurzame groei op de lange termijn. Dat sluit weer aan bij het onderzoeks-lab Analytics for a Better World dat Dick den Hertog heeft opgezet. Dit lab is op zoek naar zowel de bedrijfskundige als sociaal-maatschappelijke toepassingen van data-analyse die, zoals de naam al doet vermoeden, de wereld kunnen beteren.'

Hoe? Als voorbeeld noemt Gromicho het verkleinen van onze ecologische voetafdruk. 'Veel oliemagnaten zijn bezig hun bedrijfsmodel opnieuw uit te vinden. Data-analyse kan helpen om de transitie van fossiele naar duurzame energie te versnellen. Als zij bijvoorbeeld willen overstappen naar het opzetten van windparken, moeten ze daarvoor over veel nieuwe kennis beschikken. Neem als voorbeeld de vraag: wat is de invloed van windturbines op elkaar en op de energieproductie? Dit soort vraagstukken kun je berekenen met behulp van analytics.'

Met de opmars van computerkracht en data kunnen we nog veel meer betekenen, stelt hij. 'Met het voedsel dat we verspillen zouden we iedereen die honger lijdt van eten kunnen voorzien. Mensen in machtsposities beheersen die stromingen, waardoor het eten vaak niet op de juiste plaats belandt. Via analytics kunnen we efficiëntere wegen vinden om schaarse middelen alsnog op de juiste plek te krijgen. Zo ontwikkelde het Zero Hunger Lab van de Universiteit van Tilburg met behulp van beschikbare data een model dat het Wereldvoedselprogramma van de VN in staat stelde de voedselvoorziening in het Midden-Oosten te verbeteren en tegelijkertijd veel geld te besparen. Op die manier kan analytics macht geven aan de machtelozen.'

Het interview staat op de website van Faculteit Economie en Bedrijfskunde van de UvA (<http://bit.ly/statoruva3ujxkO>).