

DE ACTUARIS



JAARGANG 32
NUMMER 4
APRIL 2025

MAGAZINE VAN HET KONINKLIJK ACTUARIEEL GENOOTSCHAP



DATA

MODERNE OORLOGSVOERING IS ONDENKBAAR ZONDER DATA
WAT DATA ONS NIET VERTELLEN
BIJ HET CBS IS IEDEREEN EEN DATA OFFICER
DE UITDAGINGEN VAN HET MODELLEREN VAN OVERSTROMINGEN
TEKORTKOMINGEN VAN MARKOWITZ' PORTFOLIO-OPTIMALISATIE



van de redactie

Is meten ook weten?
door Elke Op het Veld – 3

interviews

De opinie van Merle Zwiers
door Paul Jurriëns – 4



Data inzetten voor betere dienstverlening
Interview met Ruben Dood
door André de Vos – 12



Data en veiligheid in de auto: is dat te belonen?
Interview met Erik Damen
door de redactie – 22

Harnessing the power of AI for actuaries
Interview Ron Richman
Jennifer Baker – 38

column

Datavolwassenheid
door Ingrid van Riel – 21

artikelen

Take one for the team!
by Rens Garssen – 8

Wat data ons niet vertellen
door Jakob Damen, Frank Pardoel en
Nathalie Cohen – 10

Framework for Financial Data Access (FIDA): brandstof voor innovatie of papieren tijger?
door Maarten Bakker en Ruben van den
Goorbergh – 16

A novel flood risk model for The Netherlands
by Radek Solnický, Arno Gabriël, Robin Nicolai and Bas Kolen – 18

Een andere kijk op data: een groter belang van datakwaliteit vraagt onderhoud van onze mentale modellen
door David Brunsveld – 24



Breakage – The magic of a loyalty programme
by Ann Cheung – 26

Deep Learning for actuaries
by Bart Custers and Garry Martes – 28

Datakwaliteit

door Pieter Bouwknecht en Anne Joosten – 30

De 'black box' uitgepakt: de mogelijke rol van Explainable Artificial Intelligence voor de verzekeringswereld
door Jens Reil, Thimo Jacobs en Axel Pison – 34

Limits to data quality?!
door Tjemme van der Meer – 36

rubrieken

Paspoorten – 32

Klimaat
Verzekeren is vooruitzien
door Daan van Ederen – 40



Scriptie
Multi-stage stochastische modellering adresseert tekortkomingen van Markowitz' portfolio-optimalisatie
door Lars Beute – 42

Pensioen
Dynamisch voorsorteren op Wtp-renteafdekking
door Ludo Nagel en Pieter Marres – 44

verenigingsnieuws

Nieuwe leden, AGenda en overig nieuws – 46
Komende thema's en colofon – 47



VAN DE REDACTIE

Is meten ook weten?

In dit themanummer staan data centraal. Vanuit verschillende invalshoeken wordt dit thema nader belicht.

Ook verzekeringen en pensioenen zijn datagedreven sectoren. Het loont dan ook om te investeren in hoogwaardige data en validatieprocessen. Data moeten correct en representatief zijn voor de werkelijkheid om te voorkomen dat verkeerde risico-inschattingen worden gemaakt. Immers: garbage in = garbage out.

De veelheid aan beschikbare data is (als deze betrouwbaar is) heel nuttig om risico-inschattingen te maken en financiële effecten af te leiden. Dat geldt zowel voor collectieve als voor individuele data.

Consumenten verwachten aan de ene kant steeds meer gepersonaliseerde informatie, maar zijn aan de andere kant vaak ook huiverig voor het 'big-brother-effect': er zijn zoveel data over ons als individu beschikbaar dat het kan voelen alsof er voortdurend iemand over je schouder meekijkt bij alles wat je doet, zeker online.

Het is voor gebruikers van data daarom belangrijk om op een ethische manier met data om te gaan. Als bepaalde data beschikbaar zijn, betekent dat dat ook dat die gebruikt moeten of mogen worden? Misschien moet je nu denken aan alle uitwassen van datamisbruik in de afgelopen jaren en is je eerste reactie: natuurlijk niet! Maar voordat je te snel reageert: ook in onze verzekerings- en pensioenwereld speelt deze kwestie. Staan we wel voldoende stil bij het ethische vraagstuk rondom de data die wij in ons werk gebruiken?

Uiteraard houden we ons allemaal aan de wetten en regels rondom beschikbaarheid en gebruik van data, databeveiliging en wat dies meer zij. Dat is een gegeven (of zou het toch moeten zijn). Daarbovenop zouden we kritisch moeten blijven op de ethische kant van datagebruik: meten is weten, maar hoever willen we daarin zelf gaan?

Elke Op het Veld
hoofdredacteur





DE OPINIE VAN MERLE ZWIERS

‘De mens moet nog kunnen overwegen’

Moderne oorlogsvoering is tegenwoordig ondenkbaar zonder data en Artificial Intelligence (AI). “We hebben te maken met digitalisering van conflicten die steeds meer op afstand worden uitgevochten”, constateert Merle Zwiars. Zij is Chief Data Officer bij het Ministerie van Defensie. Een gesprek over de toekomst, maar ook over de morele kant van war by data. “Hoe kunnen we ethisch én efficiënt reageren op de niet-ethische AI van de tegenstander?”

De afspraak is bij de Utrechtse Kromhoutkazerne, waar het innovatiecentrum van Defensie zit en waar Merle Zwiars met enige regelmaat te vinden is. Nou ja kazerne, een niet operationele kazerne, in feite een kantorenterrein. Met als verschil dat het omheind is, buitenstaanders het alleen gelegitimeerd en op afspraak kunnen betreden en je wandelend naar die afspraak tientallen militairen tegenkomt, in blauw dan wel groen tenue.

Veel jonge mannen en vrouwen die de kans lopen te zijner tijd te worden opgeroepen om in den vreemde vredestaken te verrichten of om volk en vaderland te verdedigen. Ter land, ter zee en in de lucht, maar steeds meer vanachter het bureau, op het computerscherm. Met data en AI in de hoofdrol.

“Defensie heeft altijd met informatie gewerkt”, vervolgt Merle Zwiars. “Op basis daarvan voer je een missie uit. Je gaat niet zomaar blind ergens op af. Het verschil met vandaag de dag is dat de technologische ontwikkelingen op dat vlak supersnel gaan. Vroeger was een persoon in staat om op basis van meerdere informatiebronnen iets te concluderen en als advies mee te geven aan de commandant.

Tegenwoordig heb je ongelooflijk veel data. Social media bijvoorbeeld, maar ook onze defensiesystemen die allerlei data verzamelen. Camera’s, radars, sensoren... De hoeveelheid informatie daaruit is voor mensen niet meer te behappen. Daar gebruiken we AI voor.”

BLAUWE LUCHT

Zwiars geeft een voorbeeld. “Moderne vaar-, voer- en vliegtuigen beschikken over camera’s die veel videobeelden verzamelen. Voor een groot deel bestaan de opnames waarschijnlijk uit blauwe lucht, zee, of andere niet relevante beelden. Het is zinloos om dan iemand urenlang die beelden te laten bekijken totdat hij iets ontdekt. Een AI-systeem kan bijvoorbeeld alle niet-relevante beelden eruit filteren en direct met beeldherkenning relevante informatie detecteren: ‘Om vijf uur en 33 minuten zien we dit’. Bijvoorbeeld bewegingen van een voertuig van een tegenstander. Daarop volgt een analyse en kan de analist bepalen of het interessant is om er iets mee te doen.

Op zee analyseren we vaarbewegingen. Als je dagelijks de scheepvaart-routes volgt, krijg je een patroon door de tijd heen. Dan kun je ook afwijkingen zien en bijvoorbeeld een drugsmokkel voorspellen. ‘Hé, opeens vaart er op maandag een vissersbootje op die route. Dat komt die dag nooit voor, misschien is dat iets’.”





Wat missen we als de krijgsmacht geen gebruik maakt van data en AI?

“Kijk naar wat er in de huidige conflicten al gebeurt met AI. Zo worden drones deels door AI aangestuurd, zonder dat iemand aan de knoppen hoeft te zitten. Dat is heel handig, omdat in het elektromagnetisch spectrum allerlei verstoringen kunnen plaatsvinden. Tegenstanders proberen elkaar te *jammen*: de communicatie te belemmeren, zodat verbindingen uitvallen. Maar als er geen verbinding meer is tussen operator en drone, kan die drone met succes zélf navigeren.

ANDERS STAAN WE DIRECT MET 1-0 ACHTER

AI wordt ook gebruikt bij het (automatisch) selecteren van doelwitten. Daar zijn we best wel terughoudend in. Kamikazedrones die bijvoorbeeld zelf op zoek gaan naar een doelwit. Dit kan bijvoorbeeld een Oekraïens voertuig zijn. Als die gevonden is, mag de drone er zelfstandig een bom op gooien. Als tegenstanders dat soort technologie inzetten, moeten wij ons daartegen kunnen verdedigen. Anders staan we direct met 1-0 achter.”

Waar gaat dit naar toe over pak-hem-beet tien jaar?

“Door de digitalisering worden conflicten steeds meer op afstand uitgevochten. Ook in het cyberdomein zie je allerlei ontwikkelingen, met aanvallen op digitale infrastructuur. Naar verwachting gaat het steeds meer die kant op. Dan zitten defensiemedewerkers achter schermen de vijand te bestoken. Dit gebeurt nu al met allerlei soorten drones, onbemande luchtvaartuigen en zeevaartuigen. Waarom zou je militairen onnodig in gevaar brengen?”

Straks zijn we vol in oorlog waarin we elkaar op afstand bestoken, elkaars drones neerhalen, waarvan de burgers in feite weinig van merken?

“Het is heel lastig om daarvan een toekomstbeeld te schetsen. Er is al sprake van hoogtechnologische oorlogsvoering, maar onder meer tanks

en artillerie blijven ook erg belangrijk. En we zien dat er in de huidige conflicten nog heel veel slachtoffers vallen. Waarschijnlijk wordt de moderne oorlogsvoering een combinatie van data en AI, met wapentuig zoals we dat kennen.”

Welke rol vervult de cloud?

“De ontwikkelingen in de buitenwereld met al die techbedrijven gaan veel sneller. Defensie moet daarop meeliften om te kunnen versnellen. Maar als het spannend wordt, als je gevoelige informatie gaat verwerken of je niet wilt dat informatie terug naar buiten gaat, dan halen we de informatie naar de eigen cloud, een afgeschermd omgeving.”

ETHISCHE KANT

Aan het gebruik van data en AI kleeft een ethische kant die de Nederlandse krijgsmacht voor een zeker dilemma zet. Ter verduidelijking noemt Zwiers als voorbeeld het Israëlische SAI-systeem Lavender. Dit AI-programma is, naar verluidt, in gebruik om potentiële doelwitten – personen – in de Gazastrook te identificeren op basis van grote hoeveelheden gegevens. Personen die verdacht worden van betrokkenheid bij de militaire vleugels van Hamas. Tijdens de eerste weken van het conflict heeft Lavender, ook naar verluidt, zo’n 37.000 Palestijnen als verdachte militanten aangemerkt. Critici uiten zorgen over onvolgende menselijk toezicht op het gebruik van Lavender. Ze wijzen op het risico van onjuiste identificatie, wat kan leiden tot burger-slachtoffers.

Zwiers: “Er moet sprake zijn van, zoals dat heet, menselijk oordeelsvermogen en controle. Je wilt altijd, voordat een wapen impact heeft, dat een mens nog beslist of iets een al dan niet legitiem doelwit is. Maar die Israëlische operator heeft per doelwit slechts twintig seconden de tijd om ‘ja’ of ‘nee’ te zeggen. De militairen zaten aan de lopende band ja, ja, nee, ja, nee, ja ja ja.... te tikken. Het is een treffend voorbeeld van wat de techniek in principe kan en hoe je ervoor moet zorgen dat dit op een verantwoorde manier gebeurt.”

Maar dit is niet verantwoord.
“Nee, dit is niet verantwoord. Daarom denken we in Nederland sterk na over hoe we dit moeten inrichten. Het is goed dat een systeem ondersteunt en zegt wat al dan niet een doelwit is. Maar wij vinden het heel belangrijk dat de mens begrijpt hoe zo’n systeem dat beredeneert. Er kan immers altijd ergens een foutje insluipen. De mens moet kunnen overwegen of een actie nog binnen het internationaal, humanitair, oorlogsrecht valt. Zijn we het hier ethisch en moreel mee eens? Op basis van het antwoord daarop kunnen we een beslissing nemen.

Met meerdere landen proberen we een gezamenlijk beeld te schetsen wat ethische AI is en daarover afspraken te maken. Tegelijkertijd kunnen we onze ogen niet sluiten voor tegenstanders die AI op een andere manier inzetten. Zoals de inzet van bewapende drones. Het is ook niet verantwoord als een zwerm drones op je afkomt en wij hebben mensen die op knopjes moeten drukken voordat we kunnen verdedigen. Hoe kunnen we dus ethisch én efficiënt reageren op niet-ethische AI van de tegenstander?”

DEFENSIE HEEFT ALTIJD MET INFORMATIE GEWERKT

Dat gaat niet.

“Daar moeten we over nadenken. We moeten sowieso altijd voldoen aan het internationaal humanitair oorlogsrecht, ook als we AI inzetten. En je wilt natuurlijk, als je reageert, geen burgerslachtoffers maken. Veel van dit soort randvoorwaarden kun je inprogrammeren. Net als de robotstofzuiger, die zichzelf niet van een trap stort. Bijvoorbeeld dat een drone stopt als deze een kindje ziet. Of dat drones niks mogen in stedelijk gebied. Zo kun je allerlei vangrails inprogrammeren.”

Wanneer zijn data een vloek?

“Op het moment dat we niet meer begrijpen waarom een systeem iets doet en dit negatieve gevolgen heeft. Dat geldt bijvoorbeeld ook voor desinformatie. Als we op allerlei social media echte berichten niet meer van nepberichten kunnen onderscheiden. Van propaganda die bedoeld is om ons te misleiden. Dan is data een vloek.”

En wanneer een zegen?

“Data en AI kunnen een gigantisch strategisch voordeel opleveren. Bovendien kunnen ze veel zogenaamde dull, dangerous en dirty taken overnemen en daarbij het werk van mensen leuker en veiliger maken.”

Heeft de Nederlandse krijgsmacht een achterstand ten opzichte van andere landen?

“Dat valt wel mee. Nederland loopt best voorop vanwege veel innovatief vermogen. De middelen zijn er nu ook om te investeren. Maar Defensie is geconfronteerd geweest met jarenlange bezuinigingen. We komen dus wél van een achterstand. We moeten een inhaalslag maken. Dat gaat best goed.

We willen naar echt multidomein optreden. Tijdens een inzet zijn alle krijgsmachtonderdelen al gewend om samen te werken. Maar voorafgaand daaraan wil je samen hetzelfde informatiebeeld hebben om te kunnen optreden. Een vliegtuig verzamelt beelden, maar een schip op zee ook. Hoe krijg je die op hetzelfde moment bij elkaar om te zorgen dat zowel de marine als de luchtmacht samen concluderen dat er bijvoorbeeld een Russisch schip op een bepaalde locatie ligt? Daarvoor heb je een geavanceerd IT-landschap nodig. Netwerken om beelden te koppelen en informatie uit te wisselen. Daar zetten we grote stappen in. Maar het is nog niet af.”

Hebben jullie actuarissen aan de slag in het AI-gebied?

“Dat zou ik zo niet weten. We krijgen wél een ander, meer wiskundig type militair. Daar zet Defensie sterk op in. De militair van de toekomst krijgt een andere opleiding, met een stukje data en AI, zodat deze weet hoe hij moderne systemen moet interpreteren. Daarnaast hebben we, ook voor de defensietop, allerlei cursussen en opleidingen opgezet, waardoor iedereen enige datageletterdheid heeft.” ■



Sinds oktober 2019 is Merle Zwiers (37) Chief Data Officer bij het Ministerie van Defensie. Daarnaast is zij docent statistiek & data science bij de daarin gespecialiseerde Haagse instelling Tridata.

Voordien was Zwiers Lead Data Scientist bij het UWW en de provincie Flevoland. Daarvoor vervulde zij wetenschappelijke functies bij diverse universiteiten. Tevens was zij vier jaar editor bij de academische International Journal of Housing Policy, dat onderzoek publiceert over huisvestingsbeleid en woningmarktkwesties wereldwijd.

Met een bachelor antropologie, een master in de sociologie en gepromoveerd als sociaalgeograaf lijkt het opvallend dat Merle Zwiers in de datawereld is terechtgekomen. “Mijn achtergrond in de sociale wetenschappen kent een zware nadruk op statistiek en data. Tijdens mijn studie ben ik daar ingerold. Ik vond het een interessante manier om onderzoek te doen. Het bleek dat ik er nog goed in was ook. Nee, op de middelbare school helemaal niet. Ik vind het interessant om met wiskunde allerlei maatschappelijke vraagstukken te analyseren. Zo heb ik voor mijn proefschrift data science technieken beproefd om grootstedelijke problematiek, zoals achterstandswijken en inkomensongelijkheid, en beleid daarop door de jaren heen, te analyseren. Tegenwoordig zijn dergelijke technieken enorm ontwikkeld. Bij Defensie is dat helemaal van een ongekend niveau. Zo hebben bijna alle voer-, vaar- en vliegtuigen camera's die de omgeving filmen. De beelden kun je automatisch analyseren. Wat voor voertuigen staan er bijvoorbeeld aan de overkant? ‘Met tachtig procent zekerheid zijn deze van Rusland’.”



Take one for the team!

Do you remember the substitution of Jasper Cillessen, goalkeeper of the Dutch national football team, during the quarterfinal of the World Championship 2014 against Costa Rica? Coach Van Gaal decided not to inform Cillessen of the plan, but after the match he agreed and was happy to take one for the team. Did I rewatch the highlights of the match again when writing this article? Yes. But the more relevant questions from an actuarial point of view are: ‘Are we willing to take one for the team when buying an insurance policy?’

R. Garssen MSc is Actuarial Manager at Oliver Wyman and Member of EIOPA's Consultative Expert Group on data use in insurance.



DATA VOLUMES FEED PERSONALISATION

The utilisation of data enables businesses to effectively identify the unique preferences and behaviours of each consumer. This trend is also evident in the insurance industry, where an increasing volume of data is employed to target new customers, and, more importantly, enhance risk classification.

The increase in available data has allowed for the inclusion of more variables in pricing models, resulting in unique quotes for each individual. For consumers, this development seems promising, as it enables rewards for low-risk behaviour and ensures that individuals are not subsidising the risks of others. However, the increased availability of data raises significant concerns regarding data privacy and ethical use, as the more data is utilised, the greater the risk of misuse or breaches. Furthermore, while personalised pricing may enhance the retention of low-risk profiles for insurers, it ultimately undermines the critical concept of pooling of risks, where the risks of a few are shared among many. The loss of solidarity not only threatens this support system but also risks leaving out individuals who may feel excluded from coverage because of their higher risk profiles. This observation is not new, nor are the concerns raised in the media (van der Vught, 2014, Cornelissen, 2020, and Isidore, 2025).

ORGANIZING SOLIDARITY

De Bruin & Karssing (2017) defined possible solutions to organise solidarity and mitigate the risk of an uninsured population:

- Governments organise solidarity:** The government implements a model where solidarity is mandatory, similar to the Dutch health insurance system, where standard coverage is compulsory for all, with minor premium differentiations among providers.
- Insurers organise solidarity:** Insurers are responsible for organizing solidarity by defining homogeneous risk groups and implementing appropriate underwriting conditions. These groups may be based on data; however, a solution must be defined for consumers in the uninsurable or extremely high premium categories.
- Consumers organise solidarity:** Consumers can also demonstrate solidarity. If low-risk individuals accept slightly higher premiums, this can create a stronger sense of solidarity between low- and high-risk individuals.

Each of these options requires time and has its limitations. Moreover, changing behaviour—whether among insurers or consumers—is not a quick and straightforward task.

MEASURING SOLIDARITY

The insurance sector has long been familiar with financial key performance indicators (KPIs). More recently, non-financial KPIs have been introduced, such as carbon emission reduction targets and gender

pay gap analyses. Introducing a non-financial KPI for solidarity helps to get a better understanding of the current premium differentiation and level of solidarity.

Not all insurance products are in immediate need for measuring solidarity. The need is low for products where solidarity is already present (health insurance) or where insurance coverage is strongly segregated (travel insurance). The need for measuring solidarity is high for insurance products where pricing is heavily influenced by personal data, for instance term life insurance, or where catastrophic losses could occur, for instance floods, storms or wildfires.

The Dutch Association of Insurers is already publishing a ‘Solidarity monitor’ on a bi-annual basis for the Dutch insurance sector based on a straw-man analysis (Verbond van Verzekeraars, 2023). The report concludes that no clear trend in decreasing solidarity has been observed (p. 5). However, for home and content insurance specifically, the report (p. 14 and 16) indicates increasing differentiation in premiums and a rise in rejections. If this trend continues, solidarity will definitely be reduced.

Initial requirements for a solidarity KPI are:

- The KPI can only be defined for homogeneous risk groups.
- The KPI should not expose any confidential information, meaning it should be designed to prevent the derivation of premium information from competitors.
- The KPI must consider the spread or deviation in premiums within each homogeneous risk group on an annual basis.
- The KPI should account for both accepted and rejected profiles.
- The KPI must be adjusted for the sum insured, which is unavoidably a differentiator in premium calculations.

An example for a solidarity metric fulfilling these requirements is:

$$Solidarity = \frac{\max(P_{adj}) - \min(P_{adj})}{\text{mean}(P_{adj})}, \text{ with } P_{adj} > 0$$

The adjusted premium (P_{adj}) refers to premiums adjusted for the sum insured. If we apply this formula on synthetic data set where a premium is linearly derived from the sum insured in combination with a risk differentiation factor based on 3 levels (high, medium and low), we get to the following results:

Solidarity metric	Risk differentiation			Premium measures			Solidarity
	Low	Medium	High	Min	Max	Mean	
Perfect solidarity	1.00	1.00	1.00	4,910	4,910	4,910	0.00
Low risk differentiation	0.95	1.00	1.10	4,664	5,401	4,935	0.15
Medium risk differentiation	0.70	1.00	1.20	3,437	5,892	4,625	0.53
High risk differentiation	0.30	1.00	1.50	1,473	7,365	4,282	1.38

The metric successfully shows that the solidarity measure moves towards perfect solidarity when risk differentiation decreases.

CONCLUSION

The increasing availability of data and the trend toward personalised insurance products significantly enhance insurers' ability to classify risk and differentiate premiums. However, this heightened differentiation poses a serious threat to the principle of solidarity within the insurance sector. As recognised by the Dutch Association of Insurers in their strategy for 2025–2027, titled ‘Together for Solidarity’, and echoed by EIOPA's establishment of a Consultative Expert Group on data use in insurance, there is an urgent need to address these challenges.

To prevent a scenario where high-risk individuals are perpetually left uninsured, insurers must take one for the team by actively measuring and promoting solidarity. This can be achieved by developing and implementing a solidarity KPI that encourages transparency and accountability in premium pricing. Additionally, insurers should engage consumers in a dialogue about the importance of collective responsibility, fostering a culture where low-risk individuals are willing to accept slightly higher premiums to support their higher-risk counterparts. As Jasper Cillessen can confirm, taking one for the team will result in success for the group.

The implications of decreasing solidarity extend beyond the insurance industry; they can affect public health, economic stability, and social cohesion. A fragmented insurance landscape could lead to increased financial strain on vulnerable populations, ultimately undermining the very foundation of mutual support that insurance is built upon. By taking proactive steps to safeguard solidarity, the insurance system remains equitable and sustainable for all, reinforcing the idea that we are indeed ‘in this together’. ■

References:

- de Bruin, T., & Karssing, E. (2017). Verzekeren, technische solidariteit en morele solidariteit. Kiezen tussen de-solidarisatie en reddings.
- Cornelissen, J. (2020). Data rukken op in verzekeringen, zorgen over discriminatie groeien. *Het Financieele Dagblad*.
- Isidore, C. (2025, January 17). California fire costs insurance premium hike. CNN.
- van der Vught, T. (2014). Menzis: Solidariteit in zorg in gevaar. AM.
- Verbond van Verzekeraars, Data Analytics Centre (2023), Solidariteitsmonitor 2023. *Verbond van Verzekeraars*.



Wat data ons niet vertellen

De kracht van context en perspectief

Een van de interessantere scènes uit Adam McKay's *The Big Short* is wanneer Danny en Porter (gespeeld door Rafe Spall en Hamish Linklater) op onderzoek gaan in een nieuwbouwwijk in Florida om te kijken hoe waarschijnlijk het is dat de bewoners hun hypotheek kunnen betalen? Al snel komen zij erachter dat het merendeel van de huizen verlaten is en de bewoners die zij spreken geen benul hebben wat voor woekerhypotheek zij eigenlijk zijn aangegaan. Na deze inspectie in Florida, realiseren Danny en Porter zich dat de huizenmarkt inderdaad op instorten staat en zij shorten de markt tegen het advies van iedereen in.

Al vertelden de wiskundige modellen en de bijbehorende ratings op de hypotheekobligaties ons dat de Amerikaanse huizenmarkt zeer solide was, bleek dit in de werkelijkheid verre van waar, met de kredietcrisis van 2008 als gevolg. Deze crisis wordt vaak bestempeld als een black swan event, iets wat niemand zag aankomen. Maar hoe slaagden sommigen, zoals Danny en Porter, erin om dit wél te voorzien?

Dit had naar ons idee twee hoofdredenen. Ten eerste namen zij de data en de bijbehorende aannames over de huizenmarkt niet simpelweg aan als waarheid. Zij gingen actief toetsten of de aannames wel klopten. Ten tweede bekeken zij de situatie vanuit meerdere perspectieven. In plaats van alleen te vertrouwen op de ratings op hun computerscherm, reisden zij naar Florida en zagen met eigen ogen wat er werkelijk speelde.

J. Damen MSc (links) is Senior Consultant; F. Pardoel FRM MSc (midden) is Co-founder; en N. Cohen MSc is manager; allen werkzaam voor ADC Consulting.



INZICHT 1: VERTROUW DATA NIET BLINDELINGS

Laten we eerst inzoomen op de eerste reden: zij accepteerden de data over de huizenmarkt niet blindelings, maar stelden de vraag wat hun data daadwerkelijk maten en welke aannames eraan ten grondslag lagen.

De data en aannames waarop we onze modellen baseren, vormen slechts een vereenvoudigde weergave van de werkelijkheid. Net zoals een kaart een landschap representeert, zonder elk detail vast te leggen, zo ook geeft data ons slechts een beperkte inkijk in de werkelijkheid. Data tonen patronen en geven richting, maar vangen nooit alle nuances waarop diezelfde data zijn gebaseerd.

Een kaart kan ons helpen bij het navigeren door de bergen, maar het vertelt ons niet waar de losse stenen op het pad liggen en of de route ons gaat bevallen. Op dezelfde manier bieden data inzichten, maar missen zij vaak de context en subtiliteiten die essentieel zijn voor een volledig begrip van de situatie.

Een concreet voorbeeld hiervan is de koopgeschiedenis van iemand die boodschappen doet in de supermarkt. De kassabon laat zien dat deze persoon een gezonde maaltijd heeft gekocht, maar wat de bon niet vertelt, is dat hij na enige twijfel besloot die zak chips toch maar terug te leggen. De database registreert alleen wat er daadwerkelijk is afgerekend – niet het overwegingsproces dat eraan voorafging.

In dit geval zullen de gevolgen niet zo groot zijn: misschien ontvangt deze persoon volgende week minder relevante bonusaanbiedingen en loopt de supermarkt wat omzet mis. In andere situaties kunnen de verschillen tussen wat data ons vertellen en wat er werkelijk speelt, wel grotere gevolgen hebben. Neem bijvoorbeeld data die verzekeraars gebruiken:

- **Schadeclaims** voor nieuwe klanten en schadeverzekeringen voor klimaat- en natuurrampen worden vaak gebaseerd op historische data. De vraag is echter of deze data wel representatief en volledig zijn. Kortom, sluiten de omstandigheden van vroeger nog aan bij de risico's van nu en de toekomst?
- **Levensverzekeringspremies** worden onder meer bepaald door (zelfgerapporteerde) gegevens over gezondheid en leefstijl, zoals

gezondheidsklachten en rookgedrag. Ziekte of een ongezonde levensstijl kan leiden tot significant hogere premies of zelfs uitsluiting van de verzekering. Een grote verzekeraar stelt in zijn gezondheidsformulier: 'Wij gaan ervan uit dat u de vragen eerlijk, juist en volledig invult'. Gezien de prikkels is de vraag in hoeverre men daar echt vanuit kan gaan.

Een mogelijke oplossing is om niet alleen data te gebruiken die snel, goedkoop en makkelijk voorhanden zijn, maar om te evalueren welke data écht nodig zijn om de vraag zo volledig mogelijk te beantwoorden. Hierbij ook de evaluatie of er dan voldoende capaciteit en budget beschikbaar is om deze dataset samen te stellen.

Door kritisch te kijken naar de onderliggende data voor het specifieke vraagstuk kan worden voorkomen dat er onjuiste of onvolledige conclusies worden getrokken.

INZICHT 2: HET LOONT DATA VANUIT MEERDERE PERSPECTIEVEN TE BEKIJKEN

De tweede reden waarom Danny en Porter in *The Big Short* de kredietcrisis zagen aankomen was dat zij hun data vanuit meerdere perspectieven bekeken, en niet alleen maar vanuit een kwantitatieve hoek.

Een kwantitatieve aanpak werkt goed voor onderwerpen als natuurkunde en scheikunde. Maar deze aanpak schiet snel tekort zodra menselijk gedrag verklaard dient te worden. Of zoals beroemd natuurkundige Neil deGrasse Tyson ooit zei: 'In science, when human behavior enters the equation, things go nonlinear. That's why physics is easy and sociology is hard'.

Christian Madsbjerg biedt ons een potentiële oplossing voor dit probleem. In zijn boek *Sensemaking* duikt hij in het probleem dat wij als samenleving te gefocust zijn geworden op data en dat wij tegenwoordig de wereld alleen nog maar door de kwantitatieve lens willen begrijpen. Hij pleit voor het toevoegen van meer kwalitatieve data wanneer menselijk gedrag deel van je data beslaat. Ga bijvoorbeeld naar de plek waar de data worden verzameld en observeer

en interview mensen die de data generen. Ofwel, precies wat Danny en Porter hebben gedaan toen zij naar Florida afreisden.

Dit brengt ons bij een fundamenteel inzicht: *datasets zijn geen objectieve, wetenschappelijke feiten, maar sociale constructies*. Ze worden gevormd door de keuzes en aannames van degenen die ze creëren, verzamelen en interpreteren. Het toevoegen van kwalitatieve data kan helpen om blinde vlekken en onbewuste vooroordelen in data te verminderen.

Daarnaast kan ook een interdisciplinaire aanpak bijdragen aan het verbreden van het perspectief. Zet bijvoorbeeld data-analisten, verzekeraars en beleidsmakers bij elkaar wanneer er een nieuw model moeten worden ontwikkeld. Een data-analist kijkt waarschijnlijk anders naar een dataset dan een beleidsmaker. Door verschillende expertises te combineren ontstaat er een breder perspectief en wordt de kans kleiner dat cruciale nuances over het hoofd worden gezien.

Ook een scenario-gebaseerde aanpak, zoals beschreven door George Cairns in het boek *Scenario Thinking*, kan helpen. Hierbij brengen organisaties meerdere mogelijke toekomsten in kaart en testen ze strategieën onder verschillende omstandigheden. Dit vermindert de afhankelijkheid van één enkel model en vergroot de veerkracht, waardoor verkeerde aannames minder snel tot grote risico's leiden.

CONCLUSIE

De datarevolutie geeft ons ongekende mogelijkheden trends te herkennen, voorspellingen te doen en complexe vraagstukken te analyseren. Maar zoals de kredietcrisis van 2008 liet zien, kan blind vertrouwen op data en aannames zonder kritisch na te denken over wat de data werkelijk representeren gevaarlijk zijn.

Net als Danny en Porter moeten we verder kijken dan de cijfers op ons scherm. Door onze aannames te toetsen en verschillende perspectieven toe te passen, kunnen we betere beslissingen nemen en blijven we ons bewust van het feit dat data niet 'de waarheid' zijn, maar slechts een interpretatie daarvan. ■

Barometer van de pensioentransitie – hoe staan we ervoor?



AG JAARCONGRES 2025

dinsdag 3 juni
Spant! Bussum

Dagvoorzitter: Maarten van Meerten

Sprekers: dr. Eddy van Hijum | drs. ing. Ger Jaarsma | prof. dr. Bas Werker | prof. dr. Fieke van der Lecq

Paneldiscussie onder leiding van drs. Daan Kleinloog AAG

met onder andere Chantal de Groot MSc AAG | drs. Lieke Werner AAG | drs. Rob Schilder AAG



RUBEN DOOD:
“BIJ ONS IS
EIGENLIJK
IEDEREEN DATA
OFFICER”

Data inzetten voor betere dienstverlening

Het Centraal Bureau voor de Statistiek werkt aan methoden om de grote hoeveelheid data nog effectiever in te zetten. Onder meer met inzet van AI en ‘synthetische data’. Bescherming van privacy blijft een kerntaak. “Hoe meer je data deelt, hoe kleiner de kans op fouten”, zegt chief data officer Ruben Dood.

“Bij ons is eigenlijk iedereen data officer.” Ruben Dood is chief data officer bij het Centraal Bureau voor de Statistiek. Een relatief nieuwe functie bij het CBS, die Dood sinds begin 2023 vervult. Hij is het contact met de buitenwereld als het over data gaat, en vooral de nieuwe ontwikkelingen op data-gebied.

Het CBS is spin in het web in het datastelsel van Nederland. Vandaar de opmerking van Dood: iedereen is er met data bezig. Talrijke overheidsinstanties verzamelen in Nederland gegevens. De basis wordt gevormd door de registers van onder meer gemeenten, Kadaster, Kamer van Koophandel, Belastingdienst en UWV. In totaal zijn er tien basisregistraties. Daarnaast zijn er nog eens tweehonderd sectorregistraties, waar het CBS uit put. De data komen samen bij het CBS, dat ook zelf data verzamelt.

“We willen in Nederland dat overheden makkelijk data kunnen delen en die vervolgens kunnen gebruiken voor beleidsdoeleinden. Wij brengen onze kennis en vaardigheden op statistisch gebied in. Dat delen is best ingewikkeld, want elke overheid heeft een grote mate van autonomie. Wij mogen gebruikmaken van de data van de Belastingdienst, kunnen ze zelfs opeisen, maar de data blijven van de Belastingdienst, niet van de overheid.”

**JE KUNT DATA ALTIJD ZO INZETTEN
DAT JE ALLEEN DE WENSELIJKE
CONCLUSIES ERUIT HAALT:
CHERRY PICKING**

Basis van de dienstverlening van het CBS is het Europese programma van verplichte statistieken. Alle EU-landen moeten die data verzamelen. Dan gaat het over zaken als werkloosheidscijfers en nationale rekeningen. In opdracht van de Nederlandse overheid verzamelt het CBS nationale gegevens, bijvoorbeeld over criminaliteit en onderwijs, maar ook gezondheidsdata. Overheidsinstanties kunnen daarnaast maatwerkopdrachten geven aan het CBS om bepaalde statistieken te maken. En dan zijn er nog een paar medewerkers die eigen onderzoek doen. “Dan gaat het doorgaans om ‘nice to know’-statistieken rond de actualiteit. Bijvoorbeeld bij de Brexit of rond 25 jaar homohuwelijk.”

De data van het CBS zijn ongelooflijk gedetailleerd. Elke Nederlander zit erin, vaak op veel verschillende manieren. Om te voorkomen dat de zeer persoonlijke gegevens herleidbaar zijn en kunnen worden

misbruikt, worden ze voor gebruik gepseudonimiseerd: BSN, naam, precieze geboortedatum en adresgegevens worden vervangen door een uniek nummer. Dat nummer is vervolgens wel hetzelfde voor dezelfde persoon in alle bestanden.

Deze zogenaamde microdata kunnen door het CBS en externe wetenschappers – onder strikte veiligheidscondities – worden gebruikt voor statistisch onderzoek. “In principe zou je elke Nederlander erin kunnen terugvinden op basis van een aantal unieke kenmerken. Ik ben zelf waarschijnlijk de enige 57-jarige afgestudeerde in geodesie die in Delft woont en in Den Haag werkt. Ik ben vermoedelijk makkelijk te vinden. Alleen is dat strikt verboden.”

Een stapje hoger zijn de geanonimiseerde statistische data. Die zijn voor iedereen toegankelijk, ook als het CBS de statistiek in opdracht maakt van derden. In deze open access data zijn alleen gegevens op groepsniveau te vinden.

Het CBS zit op een enorme hoeveelheid persoonlijke data over individuele Nederlanders. Hoe kwetsbaar is dat?

“Privacybescherming is onze kerntaak. Daarom worden individuele persoonsgegevens bij binnenkomst gelijk vervangen door een nummer. We zijn heel streng op hoe de microdata worden gebruikt. Wetenschappers mogen ze gebruiken voor onderzoek, maar op dat moment zijn ze al niet meer herleidbaar tot personen. Als wij geen goed beeld hebben van de bedoeling van het onderzoek en hoe er met de data wordt omgegaan, dan gaat het niet door. En de resultaten worden altijd door ons gecontroleerd. Is het veilige, niet-herleidbare statistiek?

Bij de statistische informatie ligt dat anders. Die mag iedereen gebruiken. De statistieken zijn betrouwbaar, maar partijen kunnen er op een oneigenlijke manier gebruik van maken. Je kunt data altijd zo inzetten dat je alleen de wenselijke conclusies eruit haalt: cherry picking. Op dat moment zullen we zeker aan de bel trekken. Bij die partij, bij de media. Maar veel meer kunnen we niet doen.”

Is het CBS met al die data geen buitengewoon interessant doelwit voor hackers?

“Daarom zijn onze data ongelooflijk goed beveiligd. Onze stelregel is dat het eenvoudiger moet zijn om bij de primaire bron de data te hacken dan bij ons. Veiligheid gaat boven alles. We werken met zwaarbeveiligde verbindingen. Maar dan nog: doordat we de data bij binnenkomst omzetten in nummers, zijn persoonlijke gegevens beschermd.”

Jullie gebruiken veel data die uit andere bestanden worden opgehaald, maar doen ook nog steeds enquêtes. Hoe betrouwbaar zijn die?

“Voor sommige statistieken zijn enquêtes onvermijdelijk. Als je wilt meten hoe tevreden mensen zijn over hun woonomgeving, of het producentenvertrouwen. Enquêtes doen we steeds minder. Omdat we steeds vaker bestaande data kunnen benutten. Dan heb je een veel preciezer en landsdekkend beeld, terwijl je bij enquêtes maar een beperkte groep bereikt.





We zijn in Nederland vrij ver in die omschakeling. Net als in de Scandinavische landen. Maar in Duitsland en België doen ze nog veel enquêtes. Daar houden ze ook nog fysieke volkstellingen waarbij enquêteurs van deur tot deur gaan. In Nederland doen we dat op basis van al geregistreerde gegevens. Een enorme kostenbesparing.”

Een groot deel van jullie dataverzameling is verplicht op basis van Europese afspraken. Toch zie je bij Europese statistieken vaak een waslijst aan voetnoten omdat cijfers niet vergelijkbaar zijn. Hoe kan dat?

“Dat wordt snel minder. Die harmonisatie gaat best snel. Europese cijfers zijn grosso modo goed vergelijkbaar. Zeker als het om de kerncijfers gaat. Werkloosheidscijfers waren lang verschillend per land, maar kennen nu dezelfde definitie. Maar er zijn nog altijd verschillen. Landen die veel data uit enquêtes halen, doen dat soms met net iets andere vragen, dus dat geeft verschillen. Veel Oost-Europese lidstaten hadden vroeger geen kadaster; vastleggen van bezit was niet nodig in communistische landen. Dat betekent wel een gat in de data.”

Vervuilde data zijn een groot probleem voor elke instantie. Hoe bewaken jullie de datakwaliteit?

“Omdat we steeds minder enquêteren, wordt de datakwaliteit hoger. De datakwaliteit wordt ook beter doordat we meer data delen. Bij de tien basisregistraties checken we bijvoorbeeld op congruentie. Als

iemand in het kadaster ergens bezit heeft, moet die persoon ook in het personenregister voorkomen. Op die manier kun je fouten eruit filteren. Die zijn beperkt. Een fractie van een promille. En dan zijn veel fouten vaak nog het gevolg van vertraging: gegevens die te laat worden doorgegeven. Ook bij de sectorregistraties die een wettelijke basis hebben, is de foutmarge erg klein.

We streven naar de overheid als ‘single point of truth’. Data die eenmaal zijn geregistreerd, hoeft je niet nogmaals op te geven. Vroeger moest je bij een verhuizing alle instanties op de hoogte brengen. Dat is nu niet meer nodig.”

Data zijn het nieuwe goud. Big Tech-ondernemingen worden er rijk van. Partijen als Google en Facebook kunnen heel veel informatie over mensen genereren. In hoeverre zijn dat concurrenten voor jullie?

“Wij zijn een gegarandeerd betrouwbare en onafhankelijke partij. Wij hebben data die mensen vaak verplicht moeten afstaan aan overheden. Data die controleerbaar zijn. Wij mogen die data ook combineren voor statistieken. Hoe we dat doen, is allemaal openbaar, dus transparant.

Dat ligt voor de Big Tech heel anders. Veel van hun data zijn niet betrouwbaar. Mensen geven op social media niet altijd een realistisch of eerlijk beeld van zichzelf. Waar dat soort partijen wel heel goed in is,

is ‘quick & dirty’-onderzoek. Als Google de vragen over verkoudheid ziet toenemen, kunnen ze een griepiepidemie voorspellen. Dat doen wij niet.”

Maar zij hebben natuurlijk ook allerlei data waar jullie niet aan kunnen komen.

“Dat zijn grotendeels data die wij niet nodig hebben. Partijen als Google zie ik niet als bedreiging. Zij kunnen natuurlijk ook gebruikmaken van onze open data. Maar dat kan iedereen.”

In hoeverre gaat AI het werk van het CBS veranderen?

“Large Language Models zijn nu de rage. Die technieken kunnen wij ook inzetten, zoals we dat ook al doen met machine learning en zelflerende modellen. We zijn nu aan het testen met AI. Kunnen we AI inzetten bij het zoeken naar de goede dataset? Ik noem maar wat: je wilt een statistiek voor bosbranden. Dan kun je op bosbranden zoeken, maar dan mis je mogelijk data. Misschien moet je ook op woud zoeken, en op vuur. AI kan je helpen dat soort verbanden sneller te vinden.

Ik verwacht vooral een grotere rol van AI bij de interpretatie van data. Denk aan het uitleggen wat er in een bepaalde grafiek eigenlijk te zien is of welke conclusies je kan trekken op basis van de statistiek, bij het vertalen van data naar handelingsperspectief. Dat is iets wat wij niet doen, maar daar gebruiken overheden en bedrijven die data natuurlijk wel voor. AI kan suggesties aandragen, slimmer informatie uit datasets halen.”

IK VERWACHT VOORAL EEN GROTERE ROL VAN AI BIJ DE INTERPRETATIE VAN DATA

Overheden en bedrijven willen in toenemende mate hun beleid baseren op data. Een goede ontwikkeling?

“Ik ben buitengewoon enthousiast over databased beleid. Het is altijd goed als je je bij beslissingen op feiten baseert. Maar je kunt niet leven van feiten alleen. Je moet een visie hebben, ideeën waar je naartoe wilt. Die keuzes zijn aan de politiek.

Rond de verkiezingen ging het ineens bij alle partijen over bestaanszekerheid. Maar daar is geen duidelijke definitie van. Dus kan elke partij daar zijn eigen draai aan geven, zijn eigen cijfers gebruiken. We zijn nu bezig om tot een definitie te komen, gebaseerd op feiten. Dan weet je waarover je praat.”

De enorme hoeveelheid data die er inmiddels is over mensen, is ook beangstigend. Die data kunnen op allerlei manieren worden geïnterpreteerd en misbruikt. Hoe voorkomen jullie dat er misbruik wordt gemaakt van jullie data?

“Wij springen in Nederland en bij het CBS buitengewoon zorgvuldig om met data. Bij microdata worden alle onderzoeksvoorstellen van externe partijen voorgelegd aan een ethische commissie. Die kijkt naar de doelstellingen van het onderzoek, hoe betrouwbaar een organisatie is, welk onderzoeksplan er is en of de data daarvoor goed genoeg zijn. Jaarlijks worden er enkele verzoeken afgewezen.

Data hebben altijd een ethische component. Wij publiceren ook cijfers over migratie en criminaliteit. Daar is in toenemende mate discussie over. Ook bij ons. Je moet voorzichtig omgaan met data van kwetsbare groepen om discriminatie te voorkomen.”

Of een algoritme gaat ermee aan de haal.

“Dat kun je voorkomen door van tevoren goed te testen.”

Welke nieuwe ontwikkelingen spelen er op jullie terrein?

“De overheid is bezig met de vraag hoe je data in kunt zetten voor betere dienstverlening aan de burger. Als je ziet dat iemand zijn baan verliest, kun je hem erop wijzen dat hij recht heeft op een uitkering. Dergelijk proactief beleid kan ook op andere vlakken. Wij hebben nu data over energieverbruik op maandbasis, maar als je van energiebedrijven ook actuele data hebt, kunnen we de energietransitie

veel beter ondersteunen. En je kunt misschien voorkomen dat mensen te veel voor hun energie betalen.

Het niet betalen van de zorgverzekering is vaak een indicatie van financiële problemen. Met die betaalgegevens en data van BKR kun je wellicht voorkomen dat mensen in forse schulden terechtkomen. Statistieken over bij welke groepen mensen dit het vaakst voorkomt zijn belangrijke hulpmiddelen voor preventie. Dan moeten we wel over meer data beschikken dan we nu hebben. Ook data van bijvoorbeeld zorgverzekeraars. Bedrijven zien de voordelen wel van het delen van die data, maar ze willen nooit de enige zijn. Dus dat moet je op sectorniveau aanpakken.

Een andere ontwikkeling is het maken van synthetische data. Daarbij maak je data na op basis van bestaande data. Het zijn geen echte data, maar je kunt ze wel gebruiken om bijvoorbeeld algoritmes te testen of statistieken te maken waarvoor je te weinig echte data hebt. Het voordeel is ook dat ze veel minder privacygevoelig zijn.”

Het CBS wordt een nog grotere datafabriek.

“Dat hoeft niet. Wij zijn rupsje-nooit-genoege niet. Met de bestaande mankracht van 1800 mensen kunnen we door inzet van technologie steeds meer doen, en steeds sneller. Dat hebben we tijdens corona bewezen. We gingen van sterftecijfers op maandbasis naar sterftecijfers op weekbasis. Maar de basis blijft altijd betrouwbaarheid, onafhankelijkheid en bescherming van privacy.” ■



Ruben Dood (1967) studeerde geodesie in Delft. Daar werkte hij ook kort als onderzoeker om vervolgens over te stappen naar Rijkswaterstaat. Sinds 2008 werkt hij voor het CBS. Hij is directeur beleidsstatistiek en sinds 2023 ook Chief Data Officer.



Framework for Financial Data Access (FIDA): brandstof voor innovatie of papieren tijger?

FIDA is nieuwe Europese (concept) wetgeving die financiële instellingen verplicht klantdata te delen met gereguleerde derde partijen onder de voorwaarde van afsprakenstelsels. Ook verzekeraars zullen aan deze wetgeving moeten gaan voldoen. Dit artikel verkent de inhoud, kansen en bedreigingen van FIDA voor de Nederlandse financiële sector, en specifiek de verzekeringssector, en blikkt vooruit op de voorbereidingen en het vervolg van het wetgevingstraject.

De Europese financiële sector evolueert als gevolg van data-gedreven proposities ontwikkeld door financiële instellingen, gedreven door veranderende klantbehoeften en toegenomen verwachtingen ten aanzien van gepersonaliseerde en frictieloze klantervaringen. Wereldwijd geeft 24 procent van de financiële instellingen aan dat data-platforms topprioriteit zijn wat betreft nieuwe investeringen in technologische infrastructuur in 2024¹. Dit overtreft de 22 procent die is gericht op investeringen in kunstmatige intelligentie (AI) en onderstreept het belang van het verbeteren van capabilities in data-management om klantrelaties te behouden en concurrerend te blijven. Bovendien pleiten Europese beleidsmakers voor een datagedreven economie met de EU-financiële sector als een belangrijke pijler van economische groei. Daarnaast streven zij naar het bevorderen van open data-uitwisseling die klanten beschermt, innovatie stimuleert en concurrentie bevordert. FIDA is een cruciaal onderdeel bij het bereiken van deze doelstellingen.

WAT IS FIDA?

FIDA vereist dat financiële instellingen in de EU realtime toegang tot financiële klantdata mogelijk maken via gestandaardiseerde technische interfaces (algemeen bekend als application programming interfaces, API's) voor gereguleerde derde partijen. Hierbij staat het klantbelang centraal, gebeurt data-uitwisseling altijd op basis van expliciete toestemming van de klant en ontwikkelt de Europese Autoriteit voor Verzekeringen en Bedrijfspensioenen (EIOPA) richtsnoeren voor data-gebruik ter bescherming van klanten. FIDA bouwt hiermee voort op de Payment Service Directive 2 (PSD2), maar kent een significant grotere reikwijdte. Waar PSD2 betrekking heeft op betaaldata, gaat FIDA over een brede scope productdata waar een klant vaak al toegang tot heeft via een polis, digitaal klantportaal of app. Deze scope omvat een brede set financiële producten (onder andere sparen, hypotheken, leningen, schadeverzekeringen, investeringen, crypto) voor alle klantsegmenten; dat wil zeggen consumenten, MKB en grootzakelijk.

DIT ZORGT ERVOOR DAT ALLE PARTIJEN DEZELFDE TAAL SPREKEN

Daarnaast verplicht FIDA een marktgestuurde ontwikkeling van afsprakenstelsels, waarin bindende afspraken en standaarden voor veilige en efficiënte data-uitwisseling worden vastgelegd. Dit zorgt er voor dat alle partijen dezelfde (technische) taal spreken en op basis van uniforme juridische en businessafspraken werken. Deze afsprakenstelsels omvatten dus technische afspraken over datastandaarden, maar ook afspraken over onder meer juridische aansprakelijkheid en financiële compensatie voor financiële instellingen voor het delen van data. Dit concept is niet nieuw, zo kent de Nederlandse betaalsector iDEAL en wordt het eHerkenning-afsprakenstelsel al 15 jaar gebruikt als basis voor inloggen door bedrijven bij publieke en private dienstverleners. Eerste marktinitiatieven voor ontwikkeling van FIDA-afsprakenstelsels duiden op een product-georiënteerde, nationale aanpak, wat betekent dat in Europa mogelijk meer dan 100 verschillende FIDA-afsprakenstelsels ontstaan.

Ir. M. Bakker (links) is Vice President bij INNOPAY, a business of Oliver Wyman.

R. van den Goorbergh MSc is Senior Consultant bij INNOPAY, a business of Oliver Wyman.



FIDA STAAT ONDER DRUK DOOR HET SIMPLIFICATIE PROGRAMMA VAN DE EUROPESE COMMISSIE, MAAR WAT IS NU DE STATUS VAN FIDA?

FIDA staat aan de start van de Europese triloogonderhandelingen, wat betekent dat FIDA nog conceptwetgeving is en in de ontwikkelfase zit. Wanneer deze onderhandelingen starten is nog onzeker, mede door de recente roep in Europa om versimpeling van wetgeving met als doel de administratieve druk voor Europese bedrijven te verminderen en hun concurrentiepositie te versterken. Deze roep is onder andere gedreven door de verkiezing van Trump en op basis van het Draghi rapport² en zorgt voor veel onduidelijkheid over de status van FIDA. Desondanks heeft de Europese Commissie FIDA opgenomen in het werkpakket voor 2025³, waardoor de triloog mogelijk dit jaar van start gaat.

WELKE KANSEN EN UITDAGINGEN BIEDT FIDA?

Het primaire doel van FIDA is om klanten te voorzien van grotere keuzevrijheid, op maat gemaakt financieel inzicht en effectieve financiële oplossingen. Bovendien is het de bedoeling om een infrastructuur te creëren voor een verscheidenheid aan marktpartijen (FinTechs én gevestigde partijen) om innovatieve proposities te ontwikkelen en Europese bedrijven te versterken. Een voorbeeld van zo'n propositie is om op basis van schadevrije jaren een gepersonaliseerd aanbod te doen voor het afsluiten van een verzekering bij een huurauto, real-time geïntegreerd in de digitale klantreis. Welke toepassingen uiteindelijk adoptie zullen realiseren is een marktvraag en is dus nog afwachten.

Naast de kansen die door een breed digitaal ecosysteem worden gecreëerd zijn er ook uitdagingen voor verzekeraars:

1. We zien dat de balans in de huidige waardeketen van verzekeraars onder druk kan komen te staan. Adviseurs kunnen makkelijk toegang krijgen tot klantdata op basis van toestemming van die klant en daar nieuwe aanbiedingen en adviesproposities omheen bouwen.
2. Ten tweede worden de implementatiekosten van FIDA gezien als een belangrijke bedreiging. PSD2-investeringen worden geschat op 7,2 mrd EUR⁴, gevreesd wordt dat de kosten voor implementatie van FIDA voor de gehele sector nog hoger zullen zijn gezien de brede productscope.
3. Nieuwe digitale toetreders of bestaande digitale spelers die online prijsvergelijken krijgen nieuwe mogelijkheden. In combinatie met adviseurs kan dit leiden tot een mogelijk verhoogd klantverloop als gevolg van switchgedrag en aanzienlijke omzetverliezen. Hoe groot en realistisch dit risico is, verschilt naar verwachting sterk per type product, klantsegment en geografische regio.

BALANS VAN FIDA-INVESTERINGEN TEGENOVER WAARDE VOOR FINANCIËLE INSTELLINGEN?

PSD2 resulteerde dus in aanzienlijke implementatiekosten voor financiële instellingen. Ten aanzien van FIDA verwachten marktpartijen dat deze kosten nog hoger uitvallen als gevolg van de reikwijdte van FIDA en het grote aantal producten en IT-systemen dat daardoor geraakt wordt. Tegelijkertijd uiten vertegenwoordigers van de sector hun zorgen over de mogelijk beperkte waarde (daadwerkelijke toepassingen die op schaal worden gerealiseerd en gebruikt door klanten) die FIDA oplevert in relatie tot deze investeringen. Gevreesd wordt voor een vergelijkbaar scenario als bij PSD2, waarin toepassingen en adoptie achterblijven en investeringen in infrastructuur niet worden terugverdiend.

FIDA GAAT IMPACT HEBBEN, WELKE STAPPEN ZIJN ER NAAR VOORBEREIDING EN VERVOLG

Op basis van conceptwetgeving moeten de eerste afsprakenstelsels 24 maanden nadat FIDA van kracht wordt zijn ontwikkeld en geïmplementeerd. Ondanks dat FIDA zich nog in de onderhandelingsfase bevindt, zijn de voorbereidingen voor afsprakenstelselontwikkeling gestart in 2024 door Nederlandse brancheorganisaties zoals het Verbond van Verzekeraars. Hiervoor genoemde redenen zijn met name behoud van strategische controle en de korte tijdslijnen voor de ontwikkeling van deze afsprakenstelsels. In afwachting van eerdergenoemde Europese ontwikkelingen zijn deze initiatieven gepauzeerd, met de mogelijkheid deze te hervatten zodra er meer duidelijkheid is over tijdslijnen.

Voor verzekeraars zelf geldt dat er, als FIDA van kracht wordt, naar verwachting interne IT- en complianceprogramma's worden opgezet om dataontsluiting mogelijk te maken. En dataontsluiting is pas het begin. Nieuwe commerciële kansen ontstaan voor bestaande en nieuwe partijen op basis van data die onder FIDA wordt ontsloten en klanten zullen op een makkelijke manier controle over hun digitale informatie kunnen uitoefenen in het ecosysteem dat ontstaat.

Wie, hoe en in welke mate deze kansen (worden) benut is nog onzeker, maar de investeringen die moeten worden gedaan zijn aanzienlijk. Dus dat FIDA een grote impact heeft op de financiële sector, daarover lijkt iedereen het eens. ■

1 – Celent Dimensions IT spending forecast model 2 – Celent (2024)

2 – A competitiveness strategy for Europe – Mario Draghi (2024)

3 – Commission work programme 2025 – Europese Commissie (2025)

4 – A study on the application and impact of Directive (EU) 2015/2366 on Payment Services (PSD2) – Europese Commissie (2023)



A novel flood risk model for The Netherlands

Even though significant parts of the Netherlands are susceptible to flooding or even lie below sea level, it is one of the safest deltas globally due to an extensive system of flood defenses. As of now, about two-thirds of levees comply with the Dutch Water Act. Despite the high level of protection against flood in The Netherlands, the financial sector requires adequate means to assess and quantify the remnant flood risk.

A probabilistic catastrophe model – if well designed – provides reliable flood risk insights on both local and national level that come forward to the needs of (re-)insurers, brokers and banks alike. Applications of such a model extend to CSRD/EU Taxonomy reporting, EIOPA regulatory standards, and financial stability stress-testing, etc. Aon Impact Forecasting and HKV Consultants have developed and released a catastrophe model to assess flood risk across the Netherlands. The model has been approved by the Institute for Environmental Studies at VU Amsterdam. It has already been deployed for most of the applications outlined above. In this article we describe the major challenges faced and design choices made during the development of this model.

FLOOD HAZARD

Using the current available Dutch flood risk data easily leads to an overestimation of the probability of primary defense failures. This is based on three assumptions that are made:

1. Failure events of levee sections are typically assumed to be independent. Hydraulic loads are interdependent though, and breaches in levees can significantly affect water flow, particularly as the discharge wave is altered. This results in lower flood probabilities down the water system. Furthermore, multiple breach events across multiple sections are usually not considered.
2. Emergency measures, such as ringing boils with sandbags to create counter-pressure and heightening the dike by placing sandbags, are not taken into account when determining the failure probabilities of flood defenses. However emergency measures can reduce the flood probability by a factor 1.5–4 (Lendering et al 2015).
3. The failure definition which is used in flood risk assessments and levee design is conservative. Generally, it only captures the initial stages of the failure process and does not encompass the entire failure mechanism. That is, several components of the failure process are excluded due to the lack of quantitative models.

The overestimation can result in undesired consequences like high costs for insurance and decreased willingness to invest in the Netherlands. Hence, Kolen and Nicolai (2024) have developed a novel method to more realistically estimate the probability of flooding in the Netherlands. It builds upon publicly available flood risk information,

and adds statistical methods and expert judgment tailored to delta areas with hydraulic interdependencies and flood defenses. The method allows for multiple breaches. The approximately 1,900 available flood scenarios for the primary defenses are expanded into five million unique flood events, which can be incorporated into catastrophe models. Consequently, it offers local flood hazard profiles at every single location in The Netherlands.

FLOOD EXPOSURE AND VULNERABILITY

From modelled exposure perspective, the following information is passed into the catastrophe model in the form of portfolio:

- The description of the exposure, both quantitative (sums insured) and qualitative, including occupancy (building type) and coverage (building, contents, business interruption). Optional details like construction material, basement presence, building height, age, and area can improve loss estimation accuracy.
- Exposure location, most often in form of coordinates or postal code, which is then compared with the flood extents stored in the catastrophe model.

Modelling motor hull portfolios includes additional challenges like fluctuating vehicle values and vehicles changing locations not only on regular basis but also due to evacuation effects in case of flood warning.

Once the specifications are set, the catastrophe model translates flood depths affecting individual risks in individual flooding scenarios into monetary loss by applying loss ratio defined by appropriate damage functions to the sum insured. The SSM2024 model described in (De

Bruijn et al. 2015) provides a wide range of damage functions for various occupancies. This set was enhanced to allow for building height- and construction material-specific curves using Impact Forecasting's methodologies. Moreover, the functions were adjusted with input from the QFLAT model (created by the University of Leuven and Probabilitas), which is fed with a significant database of validated loss claims including the ones after the July 2021 floods in Western Europe.

OUTPUTS AND RESULTS

The main output of a catastrophe model is a table detailing flood events (each specified by a flood depth map, breach location and probability) and corresponding loss. For any single local exposure it is straightforward to browse through the list and extract a flood risk profile, describing the exceedance probability of flood depths at the location. Subsequently, once flood depths are translated into monetary loss, the loss profile (PML) as well as annual average loss (AAL) can be obtained. Additionally, mapping tools can provide quick access to model results, offering visual overviews and risk specifications for selected locations.

The significant added value of our catastrophe model is that it can calculate AALs and PMLs for whole portfolios of exposures, not only single locations, reflecting on their geographical proximity and potential concurrent flooding as specified by the flood events.

Figure 1 shows the flooding probabilities of each location in The Netherlands resulting from the model. The local probability of flooding is highest in some parts of the river area. Along the coastal areas with very large and high dunes, the flood probabilities are smallest.

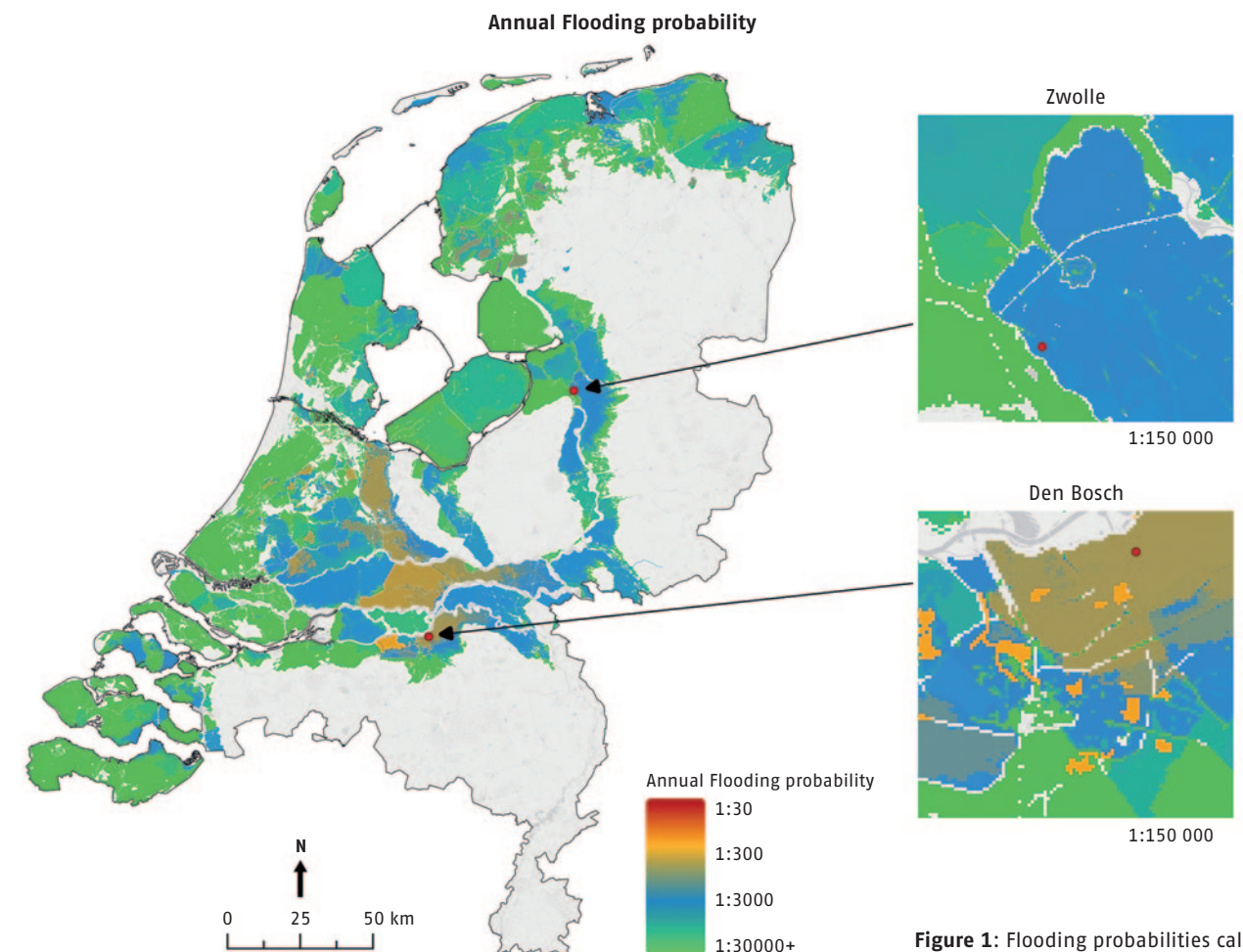


Figure 1: Flooding probabilities calculated from the model. The grey areas are higher grounds, which are not susceptible to flood risk due to primary dyke failure. Source: Impact Forecasting

f.l.t.r.:

Dr. R. Nicolai is senior advisor safety and Crisis Management at HKV.
R. Solnicki MSc is lead statistician at Aon Impact Forecasting.
Prof. dr. ir. Bas Kolen is director of research and development at HKV, senior advisor safety and crisis management at HKV and professor of Enterprise Risk Management at the UVA (part-time).
A. Gabriel MSc is managing director Capital and Risk at Aon.



Charts in Figure 2 compare flood loss profiles for two selected locations: one in Den Bosch and one in Zwolle. The annual probability of flooding for these locations is similar; it is just below 1 / 1,000. However, Den Bosch has a much higher probability of water depth exceeding 3 m. Loss profiles correspond to flood hazard profiles, but the monetary effects of flood depth increments are non-linear, yet the probability of observing loss matches the probability of flooding. The modelled building AAL for the two locations are EUR 134 and EUR 19 assuming the value of the building to be EUR 500,000.

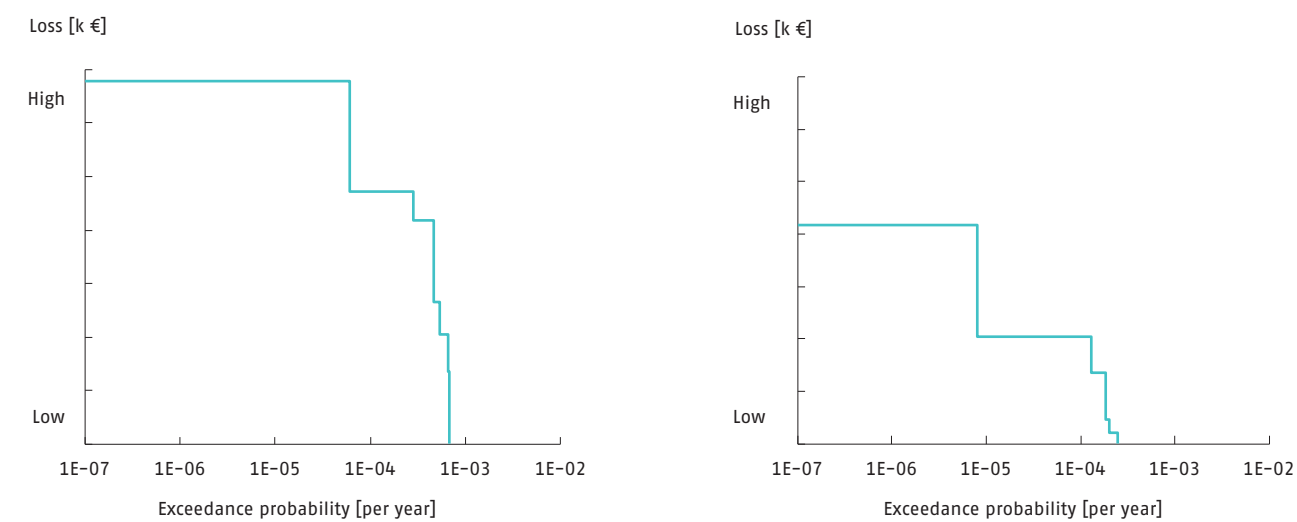


Figure 2: Flood loss profiles for two locations in The Netherlands: Den Bosch (left) and Zwolle (right).

CONCLUSIONS

Modelling flood risk in The Netherlands is a challenging endeavor. Although hydrology and defenses by themselves are well documented, the systemic risk mostly is not. Because of the difficulty to assess correlations, overestimation is a common pitfall. To correct for overestimation we have adjusted publicly available failure probabilities of flood defenses using expert judgment. Our model assesses flood damages by combining hazard, vulnerability, and exposure, producing hazard profiles, loss profiles, annual average losses (AAL), and probable maximum loss (PML) for various locations and for portfolios. Insurers are often interested in the 1 in 200 years loss. However, flood risk in The Netherlands is characterized by ‘very low probability – high impact’ events. Our model can be used to explore the effects of such events. Also, it can give input to the discussion how the Dutch society should divide flood risk between private and public parties. ■

References

de Bruijn et al. (2015)
Updated and improved method for flood damage assessment: SSM2015 (version 2). Karin de Bruijn, Dennis Wagenaar, Kymo Slager, Mark de Bel and Andreas Burzel. Deltares, 2015.

Kolen and Nicolai (2024)
A new framework to generate event sets for risk based spatial planning and financial decisions regarding Dutch Flood Management. Submitted for publication. Dutch version: https://www.hkv.nl/wp-content/uploads/2024/05/Overstromingsrisicomodel-voor-ruimtelijke-investeringsvraagstukken_R00917.10.pdf

Lendering et al. (2016)
Effectiveness of emergency measures. J. Flood Risk Manage, 9: 320–334. Lendering, K.T., Jonkman, S.N. and Kok, M. (2016). doi.org/10.1111/jfr3.12185



COLUMN

Datavolwassenheid



Datagedreven besluitvorming speelt een rol in vrijwel alles wat we doen. Een alledaags voorbeeld hiervan is boodschappen doen. Sta je in de supermarkt voor het schap (of scroll je door de app), dan beslis je of je een product koopt.

Afhankelijk van de kennis die je als consument hebt, bijvoorbeeld over de prijs, besluit je het product wel of niet te nemen.

Datavolwassenheid (data maturity) beschrijft in hoeverre een organisatie, of in dit voorbeeld de consument, data effectief en efficiënt benut. Het gaat niet alleen om de beschikbaarheid van data, maar vooral ook welke informatie je uit deze data weet te halen. Daarnaast spelen aspecten als datakwaliteit, databeheer en het gebruik van data bij strategie- en besluitvorming een rol. Er zijn verschillende niveaus van datavolwassenheid, waarvan het aantal en de definitie variëren per data maturity model. Op internet zijn diverse voorbeelden van data maturity modellen te vinden.

Om het eenvoudig te houden, illustreer ik hieronder het voorbeeld van de consument in de supermarkt op drie niveaus van datavolwassenheid: laag, matig en hoog. Hierbij focus ik alleen op de factor prijs, hoewel natuurlijk ook andere factoren van invloed zijn op het al dan niet kopen van een product.

Als de consument alleen kijkt naar de prijzen op dat moment, is sprake van een lage datavolwassenheid.

Je spreekt van een middelmatige datavolwassenheid wanneer je niet alleen de huidige prijzen bekijkt, maar ook historische prijzen en die van concurrenten. Welke prijsontwikkeling zie je door de tijd heen? Is het product eerst duurder gemaakt om vervolgens een korting te kunnen geven? Wat kost het product bij een andere supermarkt? Door deze data te verzamelen en te analyseren, kun je beter beoordelen of de huidige prijs daadwerkelijk gunstig is.

Wanneer je daarnaast gebruikmaakt van voorspellingen over de toekomstige prijsontwikkeling, spreek je van een hoge datavolwassenheid. Dit vereist niet alleen historische data, maar afhankelijk van het type product ook factoren zoals de weersverwachting en aankomende feestdagen. Door een voorspelmodel te ontwikkelen kan je beter een inschatting maken van de toekomstige prijsontwikkeling van het product. Zodoende ben je beter in staat te beoordelen of je er verstandig aan doet het product nu te kopen of dat je beter kan wachten (wanneer je een prijsdaling verwacht).

Een hoge datavolwassenheid bevordert betere, datagedreven besluitvorming. Hier ligt volgens mij een belangrijke rol voor ons als actuarieel professionals en andere data-experts: waardevolle inzichten uit data zichtbaar maken en laten zien hoe deze kunnen bijdragen aan strategie- en besluitvorming. Wat voor ons vanzelfsprekend is, geldt niet altijd voor anderen.

Ir. drs. Ingrid van Riel AAG
bestuurslid
ingrid.vanriel@actuarieelgenootschap.nl



Data en veiligheid in de auto: is dat te belonen?

SafeDrivePod is een Nederlands datagedreven technologiebedrijf met een sociale missie: het bevorderen van verkeersveiligheid door het verminderen van afleiding door smartphones tijdens het rijden en het verbeteren van rijgedrag door daarover feedback te geven. Het is in 2016 opgericht door Erik Damen en Paul Hendriks. Het idee ontstond tussen Erik en Paul tijdens hun wekelijkse fietstochten op zondag, waar ze spraken over verkeersveiligheid en de afleiding veroorzaakt door smartphones. Dit leidde tot de ontwikkeling van een technologische oplossing die het gebruik van smartphones tijdens het rijden beperkt, met als doel het aantal verkeersongelukken door afleiding te verminderen.

De redactie stelde Erik Damen enkele vragen.



E. Damen PhD is medeoprichter van SafeDrivePod en verantwoordelijk voor de technische ontwikkeling van de producten en diensten van het bedrijf. Hij heeft een achtergrond in natuurkunde en werkte vijf jaar bij de Philips Research Labs, gevolgd door vier jaar bij Philips Domestic Appliances and Personal Care. Vanaf 2000 is hij ondernemer.

Hoeveel procent van de verkeersongevallen wordt veroorzaakt door afleiding door smartphones, en hoe kan deze app helpen om dit te verminderen?

“De schattingen daarover lopen wat uiteen, omdat niet altijd bij een ongeluk duidelijk is of een bestuurder was afgeleid door de smartphone. Dat wordt niet graag toegegeven in verband met de schuld-vraag. Schattingen zijn dat ongeveer 25 procent van de ongevallen door afleiding door de smartphone wordt veroorzaakt en dat 65-75 procent van de weggebruikers toegeeft regelmatig tijdens het rijden met het telefoonscherm bezig te zijn.”

Kan de app realtime detecteren wanneer een gebruiker zich daadwerkelijk in een rijdende auto bevindt, en hoe wordt voorkomen dat passagiers onterecht worden geblokkeerd?

“De SafeDrivePod is een product dat rijstijl meet en afleiding door de smartphone voorkomt. Het is niet alleen een app, maar vooral ook een klein apparaatje (de ‘pod’) dat je achter je voorruit plakt. De bewegingssensoren in het apparaatje meten of je stil staat of aan het rijden bent en de pod zal dan het telefoonscherm sluiten. Bij de installatie van de app koppel je de geleverde pod aan de app. Dus de pod koppelt niet aan een willekeurige telefoon. Meestal is het de bestuurder van de auto die zijn of haar telefoon aan de pod heeft gekoppeld. Maar als er meerdere bestuurders van de auto zijn en de passagier ook zijn of haar telefoon aan dezelfde pod heeft gekoppeld, kan in plaats van de telefoon van de bestuurder, de telefoon van de passagier verbonden zijn. Het is natuurlijk irritant als het scherm van de passagier steeds sluit bij het rijden. Het systeem kan niet detecteren of je een bestuurder of passagier bent. Een simpele oplossing is echter dat de passagier bluetooth even uit zet. De pod koppelt dan automatisch aan de telefoon van de bestuurder en de passagier kan weer met de telefoon aan de slag.”

OF EEN BESTUURDER VAAK HARD REMT IS EEN BETERE VOORSPELLER VOOR RISICO DAN TE HARD RIJDEN

Welke soorten data verzamelt de app om rijgedrag en telefoonactiviteit te analyseren, en hoe wordt deze informatie gebruikt om veiliger rijgedrag te stimuleren?

“Er zijn meerdere meetbare parameters rond rijgedrag die het risico op schade kunnen beïnvloeden. Voor de hand ligt te hard rijden. Dit is niet eens zo heel eenvoudig te bepalen: daarvoor moet je de snelheid van de auto op ieder moment meten (GPS), maar ook de precieze locatie (ook door GPS) om te kunnen achterhalen wat de maximum snelheid op die plek is (goede kaart). Alleen dan kun je vaststellen of je ook daadwerkelijk te hard rijdt. Bij SafeDrivePod vinden we het continu volgen van een bestuurder alleen om te bepalen of deze te hard rijdt onwenselijk in verband met de privacy. En misschien wel een nog belangrijkere reden is dat, wellicht tegen de verwachting in, te hard rijden niet eens de belangrijkste parameter is die correleert met risico op schade. Een belangrijkere parameter is rembedrag. Of een bestuurder vaak hard remt is een betere voorspeller voor risico dan te hard rijden. De sensor in de pod meet de versnelling van de auto in voorwaartse, achterwaartse en zijwaartse richtingen en bepaalt daaruit het aantal matige en harde remacties, optrekken en bochten. Uit die

waarnemingen is dan een rijscore te berekenen die de bestuurder aan het eind van de rit te zien krijgt in de app.”

Kan de app gekoppeld worden aan andere rijveiligheidsoplossingen, zoals rijhulpsystemen of connected car-technologie?

“De pod en app zijn volledig zelfstandig werkende producten en niet verbonden aan de auto en dus ook niet aan rijhulpsystemen en dergelijke.”

Hoe kan de app helpen om risicogedrag, zoals herhaald telefoongebruik tijdens het rijden, te signaleren en voorkomen?

“Eigenlijk signaleren we het telefoongebruik niet, maar maken we het onmogelijk, behalve dan het toegestane handsfree bellen. Voorkomen is beter dan genezen. Om achteraf vast te stellen dat er sprake was van afleiding is interessant, maar het is beter om het ongeluk niet te laten gebeuren door het scherm te blokkeren.”

DE GENOEMDE RIJSCORE WORDT VERTAALD IN EEN KORTING OP DE PREMIE

Kunnen verzekeraars de data uit deze app gebruiken om rijgedrag te analyseren en op basis daarvan kortingen of beloningen aanbieden aan veilige bestuurders?

“Dat is precies wat we doen samen met FBTO (Achmea). De hierboven genoemde rijscore wordt door FBTO vertaald in een korting op de premie die maandelijks wordt uitgekeerd en kan oplopen tot 30 procent.”

Hoe wordt de privacy van gebruikers gewaarborgd als verzekeraars de data willen gebruiken voor risicobeoordeling?

“Zoals gezegd, gebruiken we geen locatiegegevens en zijn daarmee al de meest privacy vriendelijke telematica-oplossing. Maar de rijgedrag-data die we verzamelen met de bewegingssensoren worden gekoppeld aan een auto en daarmee aan een polis. Om de beloning te kunnen uitkeren is het natuurlijk nodig dat de verzekeraar weet over wie het gaat. Dat betekent meestal dat de data uit de app rechtstreeks naar de computers van de verzekeraar wordt gestuurd en er geen data naar SafeDrivePod gaat en wordt opgeslagen.”

Zou een verzekeraar deze app kunnen verplichten voor risicovolle bestuurders, bijvoorbeeld jongeren of bestuurders met een hoog schadeverleden, om hun premie te verlagen?

“Ook dat gebeurt nu al bij bijvoorbeeld FBTO.”

Welke impact zou het gebruik van deze app kunnen hebben op het aantal schademeldingen en de hoogte van verzekeringspremies?

“Helaas krijgen wij als SafeDrivePod geen inzage in de schades en kunnen we dus niet zelf één op één de relatie leggen tussen rijgedrag en schade. Die analyse wordt wel gedaan bij de verzekeraars. Daar komt niet een eenduidig antwoord op, omdat de doelgroepen nogal kunnen verschillen. Wij hopen in de komende jaren hier een beter inzicht in te krijgen door onze nauwe samenwerking met verzekeraars. Een indicatie voor de vermindering kan worden afgeleid uit de premie-kortingen die verzekeraars geven: in het geval van FBTO is dat dus tot 30 procent. Meestal zie je bij verzekeraars met telematica kortingen tussen de 10 procent en 25 procent.”

Kunnen verzekeraars op basis van de data uit deze app bepalen of een bestuurder afgeleid was op het moment van een ongeval, en wat betekent dat voor schadeclaims?

Als ons product op de juiste manier wordt gebruikt, dus de app op de telefoon, bluetooth aan, pod gekoppeld, dan is afleiding onmogelijk omdat het telefoonscherm dicht is. Wij kunnen monitoren of aan bovenstaande drie eisen wordt voldaan. Als dus op het moment van een ongeluk bluetooth aan stond, de app draaide en de pod was gekoppeld, dan mogen we aannemen dat de telefoon op slot was en het ongeluk dus niet was veroorzaakt door afleiding. En omgekeerd, als niet aan die voorwaarden was voldaan, dan kan afleiding niet worden uitgesloten.”



Bij de verzekeraar(s) die SafeDrivePod verplicht stellen voor risicovolle bestuurders, wat is het gemeten effect op de schadelast voor deze specifieke groep?

“Dat zouden we graag willen weten, maar die informatie hebben wij (nog) niet. De data over schadelast ligt bij de verzekeraar en niet bij ons. De analyse van het effect duurt minimaal 1-1.5 jaar, omdat er nu eenmaal echte schades moeten ontstaan en voldoende statistiek om een conclusie te kunnen trekken. De introductie van ons product bij Achmea was mei vorig jaar, dus dat duurt nog even. En daarna is het de vraag in hoeverre wij als leverancier inzage krijgen in die gegevens.”

Hoe gaat SafeDrivePod om met selectief aanzetten van de app? Mensen kunnen als ze haast hebben en harder willen rijden gewoon bluetooth uitzetten. Voeren jullie controles uit op verwachte aantal rijbewegingen versus gerealiseerde?

“We meten geen snelheidsovertredingen, want daarvoor moet je plaats en snelheid weten en wat de maximumsnelheid is, dus dat mag geen reden zijn om de bluetooth uit te zetten. Verder helpt het je niet veel: de pod slaat alle ritten op. Dus de volgende keer dat je hem koppelt komen de ritten alsnog in de lijst. Als je nooit koppelt, krijg je ook geen korting, dus dat is een goede motivator.”

OVER HET ALGEMEEN ZORGEN DE SLECHTSTE 10 PROCENT VAN DE MENSEN VOOR DE CLAIMS

Is het gebruiken van de SafeDrivePod an sich al een premieverlagende actie? Als een soort motivatie voor de polishouder het te gebruiken?

“Bij het beloningsmodel van FBTO krijg je altijd korting, behalve als je het heel bont maakt: dus dan is het gebruik an sich al premieverlagend. Bij andere modellen, die we nu aan het ontwikkelen zijn met andere verzekeringsmaatschappijen, kan dat anders werken: dan kun je bijvoorbeeld ook een hogere premie krijgen als je slecht rijdt. Je kunt per maand belonen, of alle data verzamelen om de premie voor het volgend jaar te bepalen. Er zijn heel veel manieren. In alle gevallen moeten er in de beloningsstructuur genoeg motiverende factoren zitten om het gebruik van de app en pod aantrekkelijk te maken.”

Uiteindelijk zullen alle goede bestuurders de pods blijven gebruiken, want dat levert voordeel op, en slechte bestuurders zullen de pods niet meer willen, want dat levert nadeel op. Hoe wordt de lange-termijnrelevantie van de SafeDrivePod dan gewaarborgd?

“Eigenlijk wil je als portfoli houder dit effect ook. Als de slechte risico's het portfolio verlaten, hou je dus een 'schone' portfolio over waarvan het risico als geheel veel beter te managen is. Over het algemeen zorgen de slechtste 10 procent van de mensen voor de claims. Aan de andere kant kun je je ook richten op de slechte risico's: maar dan slecht naar de klassieke maatstaven: bijvoorbeeld jonge bestuurders. Die zou je automatisch een veel hogere premie aanbieden, maar kun je een lagere premie laten 'verdiene n'.” ■



Een andere kijk op data: een groter belang van datakwaliteit vraagt onderhoud van onze mentale modellen

Wie houdt er niet van lekker eten?! Bij een goed restaurant komen we graag terug. Als culinaire fijnproever en hobbykok kwam ik erachter dat de kwaliteit van gerechten voor het grootste deel wordt bepaald door de kwaliteit van de ingrediënten. Het kookproces is ook belangrijk, maar met slechte ingrediënten wordt het nooit echt lekker. En met topingrediënten kan het bijna niet misgaan. De kwaliteit van de ingrediënten is fundamenteel. Inmiddels weet ik waar ik welke ingrediënten wil halen, kook ik met nog meer respect voor de ingrediënten, en krijg ik regelmatig complimenten.

Ingrediënten en data hebben veel gemeen. Data zijn fundamenteel voor een bedrijf, voor de maatschappij, voor ons. De digitalisering van de afgelopen decennia heeft enorme hoeveelheden digitale data gecreëerd, en zal door kunstmatige intelligentie (exponentieel) verder groeien. Digitale data zijn tegenwoordig overal en spelen een grote rol in bedrijfsstrategieën en persoonlijke beslissingen. Ik noem bewust digitale data, want er zijn ook niet-digitale data. Die hebben meer impact, maar daar zijn we ons misschien minder van bewust. Ik kom daar nog op terug.

Voor datagedreven bedrijfsvoering zijn data van voldoende kwaliteit uiteraard essentieel. Problemen door datakwaliteit kunnen zich uiten op verschillende plekken; denk aan verkeerde beeldvorming of besluitvorming door misinformatie, uitsluiting of discriminatie door algoritmen, uitval van diensten, verlies aan klanten of zelfs reputatie. Voordelen van passende datakwaliteit zijn er ook; snellere processen, lagere kosten, meer toegevoegde waarde. Actuarissen zijn grote data-gebruikers, en wij digitaliseren mee, dus alle reden dat wij ook ons begrip van datakwaliteit verder ontwikkelen.

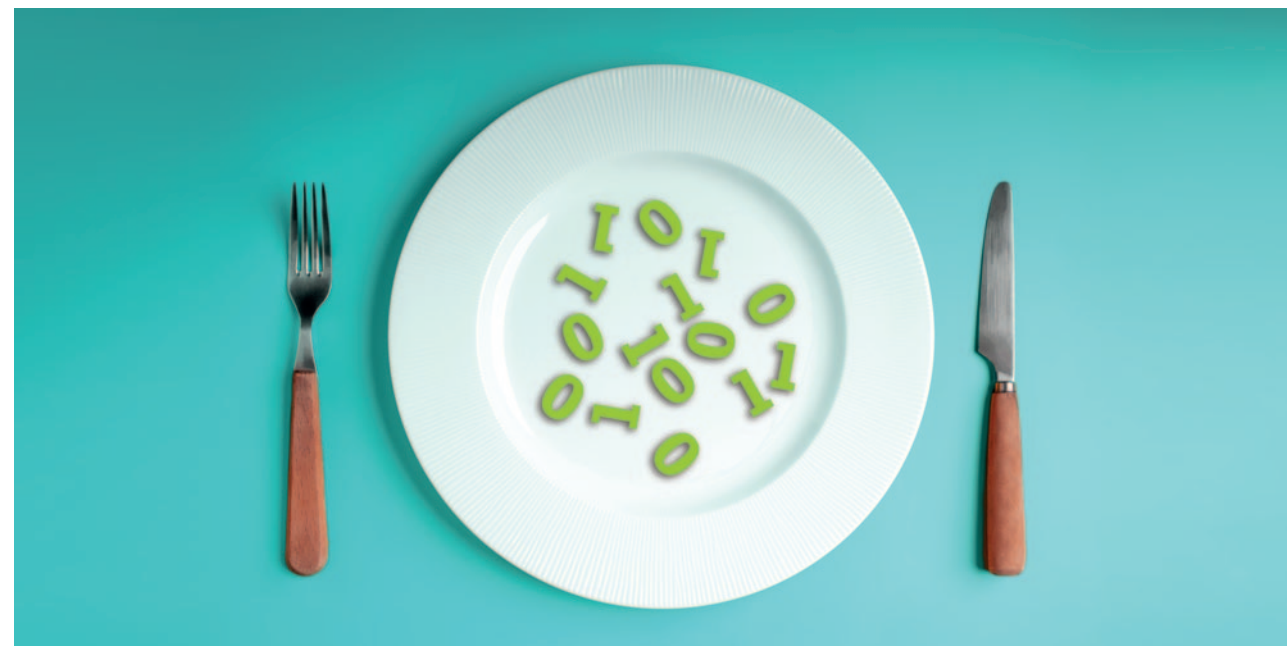
DATA MOETEN FIT FOR PURPOSE ZIJN

WAT IS DATA EN DATAKWALITEIT?

Data zijn stukjes informatie die stromen vanuit een bron naar een einddoel (vaak via IT-systemen en/of modellen), zoals naar een rapportage, een prijs, een beslissing. Datakwaliteit is een concept binnen een van de algemene standaarden voor databeheersingskaders DMBOK¹ (Data Management Body of Knowledge).

In onze praktijk wordt datakwaliteit vaak benoemd in een data-beheersingskader en in een daarvan afgeleid datakwaliteitsbeleid. Daarin staat bijvoorbeeld dat men risico's wil beheersen en dat dit betekent dat men data nodig heeft van voldoende kwaliteit; dat de data 'fit for purpose' moeten zijn voor de gebruikers. Het is 'fit for purpose' als het voldoet aan de kwaliteitsvereisten van de gebruikers. De datakwaliteit wordt uitgedrukt in dimensies (zoals compleetheid, correctheid, tijdigheid, consistentie) en gemeten op in ieder geval de kritieke data-elementen.

Onze chef wil elke klant een fijne ervaring geven. Om dat mogelijk te maken zijn alle ingekochte ingrediënten nodig (compleetheid) en geen rotte ingrediënten (correctheid) die beschikbaar zijn wanneer het proces daarom vraagt (tijdigheid). Als daarin iets misgaat, wordt dat snel duidelijk. Onze chef overziet het geheel, kiest de leveranciers, stelt hen vragen, geeft feedback, beoordeelt de kwaliteit van geleverde producten, proeft ze, en vraagt om en ontvangt feedback van gasten. De analogie, zorgen voor en onderhouden van datakwaliteit is



mensenwerk en vraagt holistische kennis. Het vraagt kennis en ervaring voor de keuze van ingrediënten, van het einddoel, en van de data-bewerkingen daartussen.

Leuk dat restaurant, maar een financiële dienstverlener is veel complexer! Er zijn alleen al zeer verschillende 'gasten'; management, polishouders, aandeelhouders, toezichhouders, de maatschappij. Onvoldoende datakwaliteit geeft ook in de financiële sector absoluut (financiële) risico's. Ook het succes en effectiviteit van een data-gedreven financiële sector hangt in belangrijke mate af van de kwaliteit van de input data; 'garbage in, garbage out'. Onze chef knikt instemmend.

ALS ER IETS MIS IS MET DATA, DAN HEBBEN WIJ DAAR SNEL MEE TE MAKEN

RISICO'S EN VOORDELEN

De toename in ons gebruik van digitale data geeft een toename in bepaalde risico's. Onjuiste, onvolledige of verouderde data kunnen leiden tot verkeerde schattingen, beslissingen of acties. En dat kan zich gaan uiten op meer plekken (bijvoorbeeld klachten van klanten, in de waardeketen). Dit benadrukt het belang van een beter begrip van datakwaliteit en hoe deze gewaarborgd kunnen worden. Actuarissen selecteren, bewerken en leveren data zoals in het stellen van aan-ames, het genereren van schattingen, pricing, et cetera. Als er iets mis is met data, dan hebben wij daar snel mee te maken.

Onvoldoende datakwaliteit geeft ook in ons actuariële werkveld (financiële) risico's, dat lijkt mij duidelijk. Dit kan bijvoorbeeld leiden tot langere doorlooptijden of hogere kosten. Denk aan benodigde analyses van waar een onverklaard verschil zit en hoe dat structureel op te lossen. Uit ervaring weet ik dat onder tijdsdruk een issue soms wordt 'opgelost' door een modelaanpassing, terwijl het echte probleem in de data zat. Onze chef zou het wel weten. Wetende dat een probleem niet aan de ingrediënten lag, richt de blik zich op het verwerkingsproces. Hij/zij zou er snel achter komen wat en waar het mis ging, het die dag nog oplossen en ervan leren. Klaar voor de volgende dag.

IS DAN EEN VERGELIJKING VAN CIJFERS VAN SOLVENCY II MET IFRS 17 TER CONTROLE NOG NODIG?

Data van goede kwaliteit geven veel voordelen. Het kan, ook voor actuarissen, tijd en kosten besparen en waarde helpen toevoegen. Stel je voor dat alle data, maar ook de modellen en IT-systemen 'fit for purpose' zijn. Is dan een vergelijking van cijfers van Solvency II met

IFRS 17 ter controle nog nodig? Dan zou er zeker meer tijd zijn voor waardetoevoegende activiteiten. Ons gebruik en belang verandert mee met digitalisering en automatisering. Het belang van datakwaliteit voor ons lijkt mij des te meer duidelijk.

HOE KIJKEN WIJ NAAR DATA KWALITEIT?

Onze kijk op datakwaliteit doet er daarom toe, nogal. Hier kom ik terug op de grote impact van niet-digitale data. Ieders blik kan verschillen, gevoed door bijvoorbeeld (ook onbewuste) aannames, overtuigingen. Deze blik wordt in systeemdynamica ons 'mentale model'² genoemd. Dit zijn extreem rijke niet-digitale databases. Onze mentale modellen helpen ons om efficiënt te beslissen, maar kunnen ons ook flink misleiden.

Ik introduceer daarom een risico: 'mental model risk'. Zeg maar: de kans op negatieve gevolgen door een mentaal model dat niet 'fit for purpose' is. In context van data betekent dit dat onze mentale modellen kunnen leiden tot verkeerde interpretaties en beslissingen. Wanneer we data interpreteren door de lens van onze mentale modellen, lopen we bijvoorbeeld het risico om alleen die informatie te zien die onze bestaande overtuigingen bevestigt, of we zien iets simpelweg niet, of we vergissen ons in wat de kritieke data-elementen zijn. Mental model risk kan zich ook uiten in een selectie van data-bronnen, waarbij onze mentale modellen ons weer leiden naar data-bronnen die onze overtuigingen bevestigen.

Dit benadrukt het belang om onze mentale modellen 'fit for purpose' te houden. De vraag is of en hoe we dat kunnen doen. Dat het kan, durf ik wel te stellen. Hoe, lijkt mij de uitdaging. Daar kunnen we misschien wat leren van model risk management. Bijvoorbeeld, door gehanteerde aannames bij besluitvorming transparant te maken en regelmatig hierover in gesprek te gaan.

CONCLUSIES?

Data zijn fundamentele bouwstenen. Door de digitalisering neemt het belang van data en data kwaliteit voor ons toe. Hoewel data op zichzelf objectief zijn, is onze interpretatie ervan dat vaak niet. Onze mentale modellen spelen een bepalende rol in hoe we data en datakwaliteit zien en begrijpen, en daarop handelen. Door ons bewust te worden van 'mental model risk' en door onze mentale modellen te onderhouden, verbeteren we de kwaliteit van onze data-interpretaties en daarmee van beslissingen. Food for thought? Eet smakelijk! ■

D. Brunsveld MSc AAG SCR heeft dit artikel op persoonlijke titel geschreven.





Breakage – The magic of a loyalty programme

A real-world example of how actuaries solve the breakage analytics problem in an airline frequent flyer programme

Breakage modelling is the primary focus of actuaries in loyalty programmes. In this context, breakage refers to the miles or points that are earned but never redeemed and eventually expire. The purpose of breakage modelling is to estimate the percentage of earned miles that will ultimately expire (this is referred to as the breakage assumption on earned miles), or to estimate the percentage of an outstanding miles balance that will ultimately expire (this is referred to as the breakage assumption on outstanding miles). As expected, the breakage assumptions, whether on earned or outstanding miles, are strongly correlated with the expiry rules of the loyalty programme, as well as their dynamic changes due to market trends and customer expectations. Therefore, it is crucial for a breakage model to be responsive to potential changes in programme rules, initiatives, and member transaction behaviours.

This article discusses the 'why', 'how', and 'what' of a simplified breakage modelling approach. It also intends to provide inspiration for the wider application of actuarial techniques.

A. Cheung FSA GMBPSS is Programme Actuary at Cathay Pacific Airways Limited.



WHY DOES BREAKAGE MATTER?

The key objective of breakage modelling is liability management. While this might sound more relevant to insurance valuation, loyalty programmes come with a liability for future redemption costs, similar to how insurance companies reserve for future claims. It involves valuing the dollar equivalent of a mile and forecasting the likelihood that the mile will ultimately be redeemed. IFRS 15 is the guiding accounting standard for the calculation of the fair value of miles and the loyalty liability. Since an earned mile will either be redeemed or expire, the following equation applies:

$$\text{Probability (Redemption)} + \text{Probability (Breakage)} = 1$$

Loyalty programmes are transforming globally to become increasingly customer-centric. One consequence of this transformation is the relaxation of miles or points expiry rules. In general, there are two types of expiry rules: time-based and activity-based. Under time-based rules, miles expire after a predefined period. Under activity-based, miles will not expire as long as there is at least one eligible transaction within a predefined period. The more relaxed the expiry rule, the lower the breakage rate.

A BRIEF INTRODUCTION TO THE METHODOLOGY

To create and validate breakage assumptions, various modelling approaches exist in the market. Inspired by these methodologies, an internal, simplified model for predicting ultimate breakage was developed. Under an activity-based expiry rule, the outstanding miles balance can be allocated into buckets based on the *time since last activity* (TSLA). The transition rate is the percentage of miles in a bucket that move to the next, while the rejuvenation rate is the percentage of miles that restart from the first bucket. In our programme, miles expire if there is no transaction activity for 18 months. Generalising from an 18-month to an x-month activity-based expiry rule is a straightforward exercise.

Considering the relationship between breakage, rejuvenation, and transition: when a mile keeps transitioning on a monthly basis from one TSLA bucket to the next, it eventually expires after transitioning through the 18th TSLA bucket. The breakage rate in the 19th month since a mile is earned is the product of all assumed transition rates for each TSLA bucket. Rejuvenation is also important because, when the miles balance in a member's account is rejuvenated due to a transaction, the miles are reset to be 18 months away from expiry. It will then take another 18 transitions for the miles to expire.

This finding tells us that, after a mile is earned – regardless of the number of rejuvenations and transitions it undergoes (in any sequence) – as long as the mile experiences a rejuvenation at any TSLA bucket followed by 18 consecutive transitions, it becomes breakage. The number of consecutive transitions before the last rejuvenation must be fewer than 18, of course. No matter how far into the future we project, the mile's journey can be described as permutations of

transitions and rejuvenations at different TSLA buckets. The only additional criterion is that there must not be more than 17 transitions followed by a rejuvenation at any time before the mile begins its final journey of one last rejuvenation followed by 18 continuous transitions, which leads to expiry. With this, breakage modelling can be simplified into decision trees of rejuvenation and transition.

THE SIMPLIFIED BREAKAGE MODEL

The input for this breakage model is the actual breakage experience by accrual month and the number of months since the miles were earned. This is calculated by dividing the actual miles expired by the total miles earned in each accrual month. The output is the expected ultimate breakage rate. The breakage model consists of two parts:

1. Reverse engineering the product of transition and rejuvenation rates at different TSLA buckets, using actual breakage rates.
2. Projecting breakage for any month since miles were earned, which can extend up to 360 months or more, equivalent to over 30 years. Over time, the breakage rate will tend to 0.

Let us define the following:

- t as a positive integer
- $T(t)$ as the transition rate at TSLA bucket t
- $R(t)$ as the rejuvenation rate at TSLA bucket t
- $B(t)$ as the breakage rate at time- t since miles earned
- $X(t) = R(t)$, where $t = 1$
- $X(t) = T(1) \times T(2) \times T(3) \times \dots \times T(t-1) \times R(t)$, where $t \geq 2$

To calculate $X(t)$:

For $t = 1$,

$$X(t) = \frac{B(t+19)}{B(19)}$$

For $t = 2$,

$$X(t) = \frac{B(t+19)}{B(19)} - \frac{B(t+18)}{B(19)} \times X(1)$$

For $t = 3$,

$$X(t) = \frac{B(t+19)}{B(19)} - \frac{B(t+18)}{B(19)} \times X(1) - X(2) \times X(1)$$

For $4 \leq t \leq 18$,

$$\begin{aligned} X(t) = & \frac{B(t+19)}{B(19)} - \frac{B(t+18)}{B(19)} \times X(1) - \left[\frac{B(t+18)}{B(19)} - \frac{B(t+17)}{B(19)} \times X(1) \right] \times X(1) \\ & - \left[\frac{B(t+17)}{B(19)} - \frac{B(t+16)}{B(19)} \times X(1) \right] \times X(2) \\ & - \left[\frac{B(t+16)}{B(19)} - \frac{B(t+15)}{B(19)} \times X(1) \right] \times X(3) - \dots \\ & - \left[\frac{B(22)}{B(19)} - \frac{B(21)}{B(19)} \times X(1) \right] \times X(t-3) - X(2) \times X(t-2) \end{aligned}$$

To calculate $B(t)$:

For $t \geq 38$,

$$B(t) = B(t-1) \times X(1) + B(t-2) \times X(2) + B(t-3) \times X(3) + \dots + B(t-17) \times X(17) + B(t-18) \times X(18)$$



The more actual breakage experience a loyalty programme accumulates, the more credible and reliable the model output becomes, just as with any other actuarial model. However, this model is particularly useful for loyalty programmes running on an activity-based expiry rules that are still in the relatively early stages of either programme launch or expiry rule change. Only $2x + 1$ months of actual normal experience is needed to provide an acceptable indication of the ultimate breakage rate, in the case of an x-month activity-based expiry rule.

Breakage modelling and assumption setting will continue to remain as one of the key priorities for actuaries working in loyalty programmes. We manage the financial liability and ensure compliance with accounting standards such as IFRS 15. The methodology discussed in this article is an example of how we – as actuaries – add value to breakage analytics problem-solving in a non-traditional field. We have attempted to understand and predict ultimate breakage rates, with an emphasis on the concept of miles rejuvenation and transition under activity-based expiry rules. As market dynamics and customer engagement are constantly evolving, actuarial techniques will always be relevant to empower strategic decision making in optimizing loyalty programme value, driving long-term sustainable success in this unique landscape.

If you are interested to discuss breakage models and possible future enhancements, feel free to reach out to ann_cheung@cathaypacific.com. ■



Deep Learning for actuaries

Recent breakthroughs in artificial intelligence (AI) have been largely driven by neural networks, particularly deep learning frameworks. These techniques have revolutionized fields such as image recognition, natural language processing, and predictive analytics. However, it can be difficult to apply neural networks in actuarial settings. First, neural networks are often considered 'black boxes' because of their complex structure with thousands of parameters. Actuaries usually produce rating structures that clearly show how premiums and risks behave over different variables, as transparency and interpretability are crucial in actuarial work. Neural networks though, only return a single prediction which makes it much more difficult to interpret model results and provide explanations to stakeholders. Second, most actuarial datasets have a tabular structure, consisting of structured numerical and categorical variables. On tabular data, traditional models such as Generalized Linear Models (GLM) or tree-based models (e.g. Random Forest, GBM) usually perform better compared to neural networks.

B. Custers MSc AAG (left) is Senior Pricing Actuary at Allianz Direct and Student MSc Artificial Intelligence.

G. Martes MSc AAG is Pricing Actuary at Allianz Direct.

This article has been written in a personal capacity.



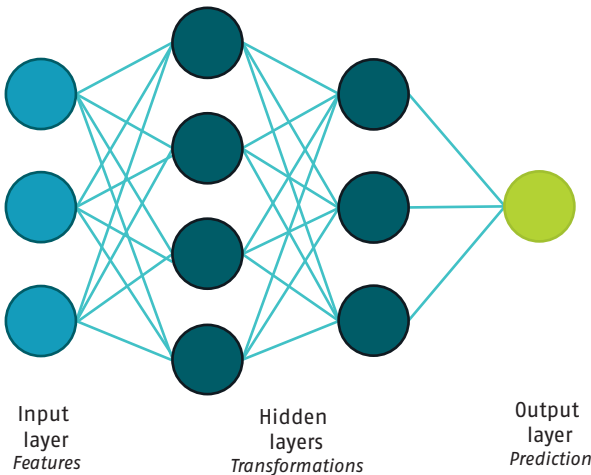
Recognizing these challenges, deep learning frameworks suited for actuarial problems and tabular data have been developed in the academic community. In this article, we discuss three of these frameworks: the Combined Actuarial Neural Network (CANN), LocalGLMNet and TabNet. We briefly explain the workings of these models as well as the pros and cons of each method.

A BRIEF INTRO TO NEURAL NETWORKS

A neural network (see figure) is a model inspired by the structure of the human brain. It consists of multiple layers of interconnected nodes (neurons) that process the data and aim to learn the relationships in the data. The three main components of a neural network are:

- **Input layer:** This is the layer that processes the input data. Each neuron in this layer represents a feature in the dataset (e.g., policyholder age, claim history, insured amount).
- **Hidden layers:** The hidden layers perform complex transformations on the data. Each neuron in a hidden layer takes inputs from the previous layer, applies a transformation and weighting, and then passes the result through an activation function.
- **Output layer:** The output layer produces the final prediction.

A typical neural network architecture



The neural network then learns by adjusting the weights of connections between the neurons to optimize the loss function and predictive accuracy.

THE COMBINED ACTUARIAL NEURAL NETWORK (CANN)

The CANN model, as proposed by Wüthrich & Merz (2019) combines a GLM with a neural network. By using the GLM predictions, the neural network part aims to capture additional patterns that were not captured by the GLM. In the output layer, the final CANN prediction is then obtained by combining the GLM prediction (that arrives via a 'skip connection') with the neural network prediction. In this way, CANN retains the interpretability of GLMs while improving predictive performance.

LOCALGLMNET

LocalGLMNet, developed by Richman & Wüthrich (2023), is another hybrid model designed for actuarial applications. Instead of producing a single set of fixed coefficients like a traditional GLM, it learns for each data record a distribution of local coefficients. For this, the model uses attention weights, which are used in a neural network to determine the importance of different parts of the input data. By multiplying the attention weights with the input data, the structure of a GLM is maintained while allowing for more complex interactions. Furthermore, the attention weights can be used to inspect variable importance and provide local explanations that can help actuaries to explain results.

TABNET

After CANN and LocalGLMNet, we can increase performance and complexity by using the TabNet model. Neural networks are known to underperform on structured data, and TabNet is a deep learning model designed specifically for tabular data. TabNet, as designed by Arik & Pfister (2021), selects the most relevant features through multiple decision steps, using a so-called attention-based feature selection. Unlike traditional neural networks that process all features simultaneously, TabNet only selects important features at each step. This is efficient and reduces the probability of overfitting. In addition, TabNet provides 'masks' that indicate at each decision step which features were found to be important and contributed to the prediction.

EVALUATION

To evaluate these models, we applied them to an open source dataset for car insurance frequency and compared performance to a GLM. The results indicate the huge potential of these methodologies. For readers who want to explore the simplified implementation in detail we provided our code and results on github.com/bart-custers/DL_for_Actuaries.

For actuaries that aim to implement these methods, performance might not be the only criterium to be considered. In the table below we aim to evaluate the deep learning techniques along four criteria: interpretability, predictive power, ease of implementation and computational efficiency. The comparison shows a classic trade-off: simpler models offer better interpretability and straightforward implementation, while more sophisticated models deliver higher performance at the cost of increased complexity and implementation challenges.

	CANN	LocalGLMNet	TabNet
Interpretability	High – Retains GLM structure	Moderate – Attention-based weighting for GLM coefficients	Low – Produces decision masks but more difficult to interpret
Predictive power	Moderate – Able to improve GLM while staying structured	High – Able to improve GLM	High – Outperforms other models in many use cases
Ease of implementation	High – Easy to implement for actuaries	Moderate – Requires hyperparameter tuning	Low – Consists of many parameters that need careful tuning
Computational efficiency	High – Simple structure, no computational issues	Moderate – Heavier in computation due to attention mechanism	Low – Can be computationally very expensive

CONCLUSION

Deep learning with neural networks has transformed many industries, but adoption in the actuarial field could be limited due to concerns about interpretability and suitability for tabular data. Models such as CANN, LocalGLMNet and TabNet aim to overcome these concerns by either combining the strengths of GLMs and neural networks or by specifically focusing on deep learning for tabular data. Our results demonstrate that these models have great potential to improve predictive performance while maintaining interpretability. Therefore, these techniques show a promising direction for actuaries that aim to leverage deep learning in their work. ■

References

Sercan Arik and Tomas Pfister. 2021. TabNet: Attentive Interpretable Tabular Learning. *In The Thirty-Fifth AAAI Conference on Artificial Intelligence (AAAI-21)*. 6679–6687.

Ronald Richman and Mario V. Wüthrich. 2023. LocalGLMnet: interpretable deep learning for tabular data. **Scandinavian Actuarial Journal** 2023(1), 71–95.

Mario V. Wüthrich and Michael Merz. 2019. EDITORIAL: YES, WE CANN! **ASTIN Bulletin** 49(1), 1–3.



Datakwaliteit

**Datakwaliteit, het begrip roept weinig passie op...
totdat je salaris naar een verkeerde rekening wordt
overgemaakt.**

**Actuariële berekeningen gebruiken complexe
modellen. Iedere actuaris weet dat het fout kan
gaan. Hoeveel berekeningen moesten worden
aangepast omdat er 'iets met de data' verkeerd was?
Een nieuw basisbestand, andere verzekerde
bedragen, een verkeerde curve: er kan veel misgaan.
Dat is (gebrek aan) datakwaliteit. In de actuariële
opleiding komen de modellen uitgebreid aan de
orde, en de actuaris vindt het leuk om modellen
beter (en complexer) te maken. Maar de modellen
worden gevoed met basisgegevens, data. Als de
basisgegevens niet kloppen, verliest de model-
uitkomst zijn waarde. Als actuaris wil je erop kunnen
vertrouwen dat de data goed zijn, maar hoe krijg je
deze zekerheid?**

Drs. P.M. Bouwneght AAG (links)
werkt bij de afdeling
Balansmanagement van
Nationale – Nederlanden Leven.

A. Joosten MSc is senior
consultant bij EY Actuarissen
B.V.



Accountants spelen een cruciale rol als partijen een oordeel willen krijgen over datakwaliteit. Dat geldt bij financiële verslaglegging, maar ook bij de overgang van pensioenfondsen naar het nieuwe pensioenstelsel. We spreken met twee accountants over datakwaliteit. Jennifer van Eekelen van EY en Wilfred Kevelam van KPMG zijn beide accountant met langjarige ervaring bij pensioenfondsen. Ze zijn betrokken bij zowel de jaarwerkcontrole als de additionele vereisten die gelden bij de overgang naar de Wtp. Het zijn de gesprekspartners die we zoeken. We gaan op zoek naar een antwoord op de vraag 'welke rol speelt de actuaris bij datakwaliteit?'. Vanwege de achtergrond van Jennifer en Wilfred gaat het vooral over pensioenfondsen, maar veel ervaringen gelden ook voor de actuaris buiten het pensioendomein.

Drs. W. Kevelam RA is Partner en Head of Pensions Audit bij KPMG.



Binnen jaarrekeningcontroles speelt datakwaliteit een belangrijke rol. Daarbij geldt een materialiteitsgrens gerelateerd aan de omvang van het fonds. Immers, de materialiteit hangt samen met het doel van het proces, de jaarrekening. Een beperkte afwijking van pensioen-aanspraken van een individuele deelnemer is voor dat doel wellicht acceptabel, maar bij pensioeningang is er nauwelijks tolerantie voor afwijkingen. Door de Wtp komt datakwaliteit onder een vergrootglas. Voor pensioenfondsen die willen 'invaren' vereist de wet¹ dat de datakwaliteit voor, tijdens en na de transitie geborgd is. Daarbij gelden twee toetsen. Bij toetsmoment 1 vraagt het fondsbestuur een auditor om specifieke werkzaamheden uit te voeren, voordat het plan om pensioenaanspraken 'in te varen' bij DNB wordt ingediend. Deze werkzaamheden hebben betrekking op de juistheid en volledigheid van de relevante pensioendata. Toetsmoment 2 geldt na de transitie: zijn alle berekeningen als onderdeel van het invaarproces correct verlopen? De opdracht aan de accountant bij toetsmoment 1 is anders dan bij een jaarwerkcontrole. De accountant kijkt of de datakwaliteit

voldoet aan de door het bestuur gestelde eisen. Het is geen zelfstandige 'assurance' op de berekening van de aanspraken. Het Kader Datakwaliteit van de Pensioenfederatie² geeft hierbij een concrete uitwerking van de wettelijke eisen.

Wilfred merkt op dat het doorlopen van het Kader Datakwaliteit veel vergt van de pensioenfondsen. Niet zozeer vanwege achterstallig onderhoud, maar met name doordat de datakwaliteitsspecificaties vrij omvangrijk zijn; alles moet namelijk gedocumenteerd worden. De nadruk ligt meer op aantoonbaarheid dan eerder.

DE NADRUK LIGT MEER OP AANTOONBAARHEID DAN EERDER

Welke ervaringen hebben Jennifer en Wilfred opgedaan? Ze verdelen de bevindingen in drie categorieën. De eerste categorie omvat *sanity checks* op de data: zijn er geen negatieve waarden, ontbreken er namen, is er iets ingevuld, et cetera. Uit deze categorie komen veel bevindingen, maar deze hebben meestal geen impact op de bepaling van de pensioenaanspraak. De tweede categorie zijn bevindingen die aan de pensioenregeling zelf zijn gelieerd, bijvoorbeeld met betrekking tot overgangsregelingen. Hier komen ook wat bevindingen uit, maar al wel een stuk minder. De laatste categorie zijn de complexe mutaties, zoals arbeidsongeschiktheid, nabestaandenpensioen en individuele waardenoverdracht. Het aantal absolute bevindingen in deze categorie is vaak relatief laag, maar deze kunnen wel impact hebben op de aanspraken. Kortom, er zijn veel punten die opgelost moeten worden, maar het heeft de invaarplannen tot nu toe niet in de weg gestaan.

Datakwaliteit verschilt per organisatie. Zijn er specifieke signalen die kunnen wijzen op dataverwaarlozing? Veel pensioenfondsen besteden de meeste activiteiten uit aan pensioenuitvoeringsorganisaties. Het niveau van datakwaliteit hangt af van de complexiteit van de pensioenregeling en de historie van het fonds met de uitvoerders. Wanneer bijvoorbeeld pensioenfondsen vaak switchen van pensioen-uitvoerder, kan dit leiden tot het verlies van gegevens, wat een nadelig effect heeft op de datakwaliteit. Bovendien hebben diverse uitvoeringsorganisaties hun eigen databewaarstrategie.

Een belangrijk moment voor de datacultuur bij pensioenfondsen was het Quinto-P datakwaliteitsonderzoek door DNB in 2011. Hierbij moesten de pensioenaanspraken van een selectie aan deelnemers met de hand worden nagerekend, wat pensioenfondsen als moment hebben aangegrepen om hun datakwaliteit te verbeteren. Het onderzoek kostte veel tijd en aandacht. Sommige pensioenfondsen gebruikten de resultaten van het onderzoek om de datakwaliteit permanent naar een hoger plan te tillen, maar andere pensioenfondsen zagen het als eenmalige exercitie, en gaven er geen structureel vervolg aan.

EEN BELANGRIJK MOMENT VOOR DE DATACULTUUR BIJ PENSIOENFONDSEN WAS HET QUINTO-P DATAKWALITEITSONDERZOEK

Zijn de eisen aan datakwaliteit veranderd door de tijd heen? Nee en ja. Nee, omdat het van alle tijden is dat iedereen recht heeft op een goed berekend pensioen. Ja, omdat technologie aanzienlijk is vooruitgegaan. Vroeger werden lijsten met mutaties handmatig ingevoerd, dat gaat nu automatisch met *straight through processing*. Je ziet dat veel werkgevers de loonadministratie hebben uitbesteed aan professionele partijen, die ook een betrouwbare aanlevering aan de pensioenuitvoerder kunnen verzorgen. Als gevolg van de Wtp stappen pensioenfondsen over naar nieuwe pensioenadministratiesystemen of maken ze de huidige systemen Wtp-proof. De nieuwe systemen bevatten meer geautomatiseerde controles, wat een positieve impuls geeft aan de datakwaliteit.

Tegenwoordig worden er experts ingeschakeld om de datakwaliteit inzichtelijk te maken door middel van tools. Een groot voordeel is dat nu de omvang van databestanden vaak niet meer voor problemen zorgt. Maar sommige dingen zijn in de loop der tijd niet veranderd: wanneer er geen historische data beschikbaar zijn, moeten pensioen-aanspraken handmatig volledig worden nagerekend. Vaak hebben actuarissen hier een rol in.

Drs. J. van Eekelen RA is Associate Partner of Financial Services Risk Consulting bij EY.



Voor de jaarwerkcontrole waren controles op basisgegevens voor-namelijk gericht op de totale voorziening. Door de introductie van het Kader Datakwaliteit is er meer focus op de datakwaliteit van de individuele deelnemer. Niet alleen de klassieke eisen aan juistheid, volledigheid en tijdigheid gelden, maar er is meer aandacht voor aantoonbaarheid. Die aantoonbaarheid leidt tot het expliciet maken van controles en normen, dat is voor veel fondsen een grote klus. Tegenwoordig kunnen veel data worden opgeslagen maar vanwege wetgeving moet er een goede afweging worden gemaakt of en hoelang data bewaard of verwijderd moeten worden.

Wilfred en Jennifer benadrukken de belangrijke rol van actuarissen in het verbeteren van datakwaliteit binnen de Wtp, bijvoorbeeld als sleutelfunctiehouder. Ze verwachten een groot piekmoment voor actuarissen bij de Wtp-overgang. Is er voldoende capaciteit?

Bij de Wtp zullen actuarissen bij datakwaliteit betrokken zijn door -heel klassiek- complexe berekeningen (handmatig) na te rekenen. Jennifer hoopt dat actuarissen vooral een rol kunnen spelen bij plausibiliteitscontroles: zijn de uitkomsten logisch? Dat vereist kwantitatief inzicht en kennis van alle onderdelen van de berekeningen: een typische klus voor actuarissen. Deze Wtp-punten gelden niet alleen voor actuarissen die zich met pensioen bezighouden, het gaat om generieke competenties van actuarissen. Daarnaast moeten actuarissen ook zelf kritisch zijn op datakwaliteit: vraag waar data vandaan komen en denk logisch na of de data geschikt zijn voor jouw doel. Ook al ligt datakwaliteit primair binnen het domein van de accountant, als actuaris heb je een rol in datakwaliteit. ■

¹ – artikel 46 lid 4 sub b van het Besluit uitvoering Pensioenwet en Wet verplichte beroepspensioenregeling

² – Pensioenfederatie, 11 oktober 2022. Het document heeft feedback van DNB, AFM, Nederlandse Beroepsorganisatie van Accountants (NBA) en diverse pensioenuitvoerders verwerkt

Actuarissen en data in hun werk



■ Welk data onderwerp staat nu bovenaan jouw to-do list?

■ Hoe is werken met data nu anders dan aan het begin van je carrière?

■ Als je een dataset zou kunnen zijn, welke zou je dan zijn en waarom?

■ Welke datafout heb je zelf wel eens gemaakt?

Naam ir. Laurens van Beurden
MSc AAG

Functie Risk Manager bij
Nationale-Nederlanden

Standaardisatie van data voor mijn werk. Door data te voorzien van duidelijke definities en deze te standaardiseren wordt de verwerking ervan vele malen makkelijker en wordt de consistentie tussen rapporten verbeterd.

Het grootste verschil zit in de verwerkingscapaciteit van computers. Waar ik vroeger nog voorzichtig moest zijn met het volume aan data, speelt dit in de meeste gevallen nu niet meer.

Ik zou niet een specifieke dataset willen zijn. Voor mij gaat het om de samenhang van data, de kwaliteit en consistentie ervan. Dus de coherente, gestandaardiseerde set aan data is waar ik voor ga.

Als natuurkundige ben ik betrokken geweest bij de automatisering van diverse meetopstellingen. Bij mijn eerste klus had ik veel assumpties over hoe de data er uit zouden zien en op basis daarvan had ik mijn programma gemaakt. Na veel geworstel bleek bij inspectie dat de data die ik ontving er door een verrassende wending niet uit zagen zoals ik had gedacht. Dus aannames zijn mooi, alleen controle van de aannames en observeren en meten van wat er echt gebeurd is noodzakelijk.



■ Welk data onderwerp staat nu bovenaan jouw to-do list?

■ Hoe is werken met data nu anders dan aan het begin van je carrière?

■ Als je een dataset zou kunnen zijn, welke zou je dan zijn en waarom?

■ Welke datafout heb je zelf wel eens gemaakt?

Naam Niek Timmerman

Functie Senior Consultant bij
EY Actuarissen B.V.

Momenteel ben ik onderdeel van een team dat bezig is met het opzetten en ontwikkelen van een datawarehouse. Dit omvat het verwerken van data afkomstig van een externe partij en het opslaan ervan op een gestructureerde manier. Hierbij komen verschillende aspecten kijken, zoals het ontwerpen van een logische datastructuur, het werken met API's, het historiseren van gegevens en het uitvoeren van transformaties. Alles wat hiermee te maken heeft, heeft momenteel mijn hoogste prioriteit.

Het werken met data is de afgelopen jaren sterk veranderd. Waar voorheen veel in Excel werd gedaan, worden tegenwoordig steeds vaker technieken zoals SQL, Python en Power BI gebruikt. Daarnaast heeft de opkomst van generatieve AI het werken met data toegankelijker gemaakt. Hierdoor zijn er meer mogelijkheden ontstaan, bijvoorbeeld doordat programmeren eenvoudiger en laagdrempeliger is geworden.

Een ongestructureerde dataset. Dat zegt genoeg, toch?

Tijdens het uitvoeren van een ETL-proces (Extract, Transform, Load) heb ik weleens ontdekt dat er een fout was gemaakt bij het overzetten van data van de bron naar de nieuwe locatie en structuur. Dit soort fouten kunnen impact hebben op de betrouwbaarheid van de data en benadrukken het belang van zorgvuldige controles en validatie.



■ Welk data onderwerp staat nu bovenaan jouw to-do list?

■ Hoe is werken met data nu anders dan aan het begin van je carrière?

■ Als je een dataset zou kunnen zijn, welke zou je dan zijn en waarom?

■ Welke datafout heb je zelf wel eens gemaakt?

Naam Joey de Mol MSc

Functie Senior Consultant Data
& AI bij Triple A Risk
Finance

Het beschikbaar maken van data bij pensioenfondsen om te voldoen aan de nieuwe pensioenwet en bovenal om deelnemers goed te kunnen begeleiden. Daarvoor ontwerp en implementeer ik moderne data-architecturen in combinatie met datakwaliteitstooling om te zorgen dat alle historische en actuele pensioendata uit verschillende databronnen consistent en betrouwbaar wordt ontsloten en daarmee snel inzetbaar is voor diverse rekentools, dashboards, en analyses.

Er is meer aandacht voor het automatiseren en professionaliseren van dataprocessen. Dat geldt voor rapportages maar ook voor machine-learningmodellen. Daarbij zie je een verschuiving naar cloudoplossingen en moderne, open source programmeertalen die geoptimaliseerd zijn voor het werken met grote hoeveelheden data. Modelleren gebeurt nu met hulp van AI en de volgende stap wordt om AI ook verantwoord in te zetten in de processen daaromheen, zoals validatie en documentatie.

De dataset van alle Strava-runs. I may trust the process, but it's all about the outcome (en kudos).

Bij mijn eerste opdracht waren er datarijen verdubbeld na een urgente wijziging met als gevolg dat er diezelfde dag nog klanten dubbel uitbetaald werden. Vanwege de directe gevolgen voor de klant is dat mij altijd bijgebleven, en adviseer ik altijd om duidelijke testprocessen en -omgevingen op te zetten.



■ Welk data onderwerp staat nu bovenaan jouw to-do list?

■ Hoe is werken met data nu anders dan aan het begin van je carrière?

■ Als je een dataset zou kunnen zijn, welke zou je dan zijn en waarom?

■ Welke datafout heb je zelf wel eens gemaakt?

Naam Siard Jongsma AAG MSc

Functie Actuaris bij ABN AMRO
Verzekeringen

Ik ben nu de jaarlijkse rapportage naar DNB aan het afronden, hieraan ligt veel data ten grondslag. Zodra alle gegevens gevuld zijn, ben je er nog niet: alle validatiechecks tussen de templates moeten ook aansluiten, tot op de cent nauwkeurig. Een puzzeltje waarbij je elk jaar probeert het aantal fouten omlaag te brengen. Inmiddels zijn veel validaties in het maakproces gebouwd zodat dit elk jaar vlotter gaat.

Data is inmiddels veel meer gecentraliseerd. Toen ik begon, was voor de meeste rapportageprocessen een aparte query gedefinieerd die een los template produceerde. Nu is het aantal databronnen sterk verminderd, wat een efficiëntere keten heeft opgeleverd. Tegelijk denk ik dat deze beweging nog verder kan worden doorgezet, dit is ook iets waar onze afdeling sterk op inzet.

Geweldige vraag. Ik kies voor de 'marathon world record progression': een dataset waarin aan elk datapunt een uniek verhaal hangt en waar je bovendien een prachtige curve doorheen kan trekken.

Los van het nulletje te veel/weinig en de verkeerde jaarlaag als input hanteren: vast meer dan waar ik weet van heb. Aangezien ik verder niet meteen iets concreets kan herinneren, zullen de consequenties daarvan tot nu toe zijn meegevallen!



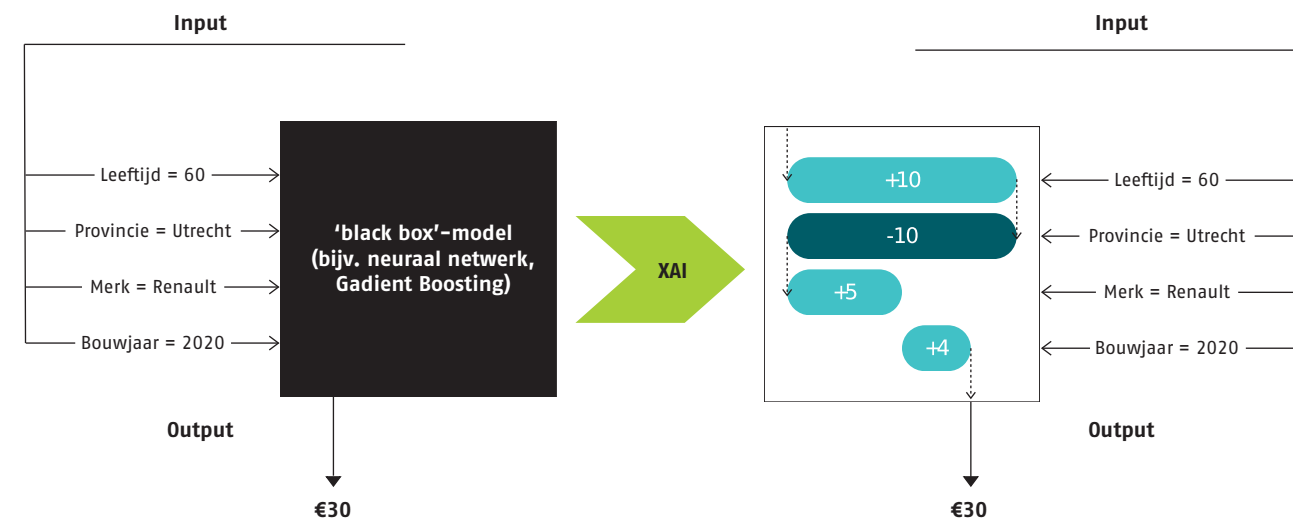
De 'black box' uitgepakt: de mogelijke rol van Explainable Artificial Intelligence voor de verzekeringswereld

Kunstmatige intelligentie (AI) is een onmisbare trend geworden, maar deze snelle ontwikkeling brengt ook diverse uitdagingen met zich mee. Om AI op een waardevolle manier in te zetten voor de verzekeringswereld is het belangrijk dat bepaalde uitdagingen worden overwonnen, zoals het 'black box'-karakter van de techniek. Dit artikel zal dieper ingaan op hoe Explainable Artificial Intelligence (XAI)-technieken de toepassing van AI binnen de verzekeringsbranche kunnen verbeteren en uitdagingen zoals 'black box'-algoritmes mogelijk kunnen oplossen.

Consumenten en toezichthouders eisen steeds vaker dat de besluitvormingsprocessen van AI transparant, duidelijk en rechtvaardig zijn. Daarnaast is de eerlijkheid van een model een belangrijke parameter die inzichtelijk moet zijn tijdens de ontwikkeling en implementatie van zo'n model. Met XAI kan men duidelijke en begrijpelijke inzichten bieden in de besluitvormingsprocessen van complexe AI-algoritmes, en opent daarmee de 'black box'-structuur van deze AI-modellen. Het uiteindelijke doel is om AI-systemen te creëren die gemakkelijk door mensen begrepen en geïnterpreteerd kunnen worden. Binnen de wereld van AI moet er op dit moment een keuze worden gemaakt tussen uitlegbaarheid of een hoog prestatieniveau. Goed uitlegbare modellen, zoals lineaire/logistische regressies, zijn vaak niet in staat om een hoge accuraatheid te bereiken. Aan de andere kant zijn neurale netwerken juist heel accuraat, maar daar is het door een grote hoeveelheid aan (intergeconnecteerde) lagen lastig om inzichtelijk te krijgen hoe de output bepaald is.

XAI is een verzameling van technieken die, in vergelijking tot neurale netwerken, uitlegbare modellen produceren én daarbij een hoog prestatieniveau behouden. Enkele voorbeelden van zulke XAI-technieken zijn LIME (Local Interpretable Model-agnostic Explanations) en SHAP (SHapley Additive exPlanations). LIME richt zich op het uitlegbaar maken van individuele voorspellingen. Het onthult hoe kleine aanpassingen aan specifieke datapunten de voorspelling van een model kunnen beïnvloeden. Dit helpt gebruikers om te begrijpen hoe een model zou kunnen presteren onder verschillende omstandigheden rondom een specifiek voorbeeld. SHAP biedt, naast het inzicht in hoe individuele voorspellingen tot stand komen, ook inzicht in de globale impact van kenmerken binnen een model. Dat maakt dat SHAP een completere XAI-techniek is dan LIME. Hoe SHAP inzicht geeft in een individuele voorspelling, is zichtbaar in het figuur. Stel dat we een 'black-box'-premiemodel hebben. Om tot een premie voor één individu te komen, wordt het premiemodel met een aantal kenmerken van een individu gevoed. Een standaard black-box model zal een premie geven, zonder dat duidelijk is hoe het model tot deze output is gekomen (linkerkant). Met behulp van SHAP kan het plaatje aan de rechterkant worden geconstrueerd. Hierin kan precies worden nagegaan welke kenmerken van een individu én in hoeverre deze bijdragen aan de output van het model. Zo zorgt in dit voorbeeld het hebben van een leeftijd van 60 voor een effect van +10 op de premie (ten opzichte van de gemiddelde output van het model).

J. Reil MSc (links) is een PhD kandidaat bij de Bern Business School, maar schreef dit artikel in dienst van Triple A – Risk Finance, waar ook consultant dr. T. Jacobs (midden) en senior consultant A. Pison MSc werken.



Figuur 1: een standaard black-box-voorspelling (links) versus een uitlegbare XAI-voorspelling (rechts)

Een andere interessante techniek binnen het XAI-domein is de zogeheten 'Explainable Boosting Machine' (EBM). EBMs zijn ontworpen om de hoge nauwkeurigheid van traditionele *boosting machines* te behouden, terwijl ze tegelijkertijd de uitlegbaarheid verbeteren. EBMs gebruiken een additief model dat bestaat uit een reeks eenvoudige, interpreteerbare componenten (beslisbomen). Dit maakt het mogelijk om de bijdrage van elke component afzonderlijk te begrijpen. Een EBM is in feite een 'glass box'-model en heeft dus geen additionele technieken zoals SHAP of LIME nodig om uitlegbaar te zijn. Door deze eigenschappen zijn EBMs interessant voor verzekeringsmaatschappijen die willen begrijpen hoe hun AI-modellen tot bepaalde beslissingen komen, zonder dat ze moeten afzien van de nauwkeurigheid van hun voorspellingen.

Met de inwerkingtreding van de Europese AI Act, die sinds 1 augustus 2024 van kracht is, is de rol van XAI steeds belangrijker geworden. De AI Act stelt namelijk eisen aan de ontwikkeling en het gebruik van AI-systemen om ervoor te zorgen dat AI veilig, betrouwbaar en transparant is. De AI-toepassingen worden geclassificeerd op basis van bepaalde risico's en de AI Act stelt daarbij verschillende eisen afhankelijk van het risiconiveau. Voor hoog-risico AI-systemen, zoals die gebruikt worden in de financiële dienstverlening, zijn duidelijke eisen gesteld om de risico's voor fundamentele rechten te minimaliseren. XAI helpt hierbij door de werking van deze systemen te verduidelijken, zodat fouten sneller kunnen worden opgespoord en gecorrigeerd. Daarnaast moeten ontwikkelaars van AI-systemen de risico's voor fundamentele rechten in kaart brengen en klanten actief informeren over de inzet van hoog-risico AI-systemen. Dit zorgt ervoor dat klanten bewust zijn van wanneer AI wordt gebruikt in beslissingen die hen direct raken, zoals bij financiële beoordelingen. Door de besluitvormingsprocessen transparanter te maken, kunnen verzekeringsmaatschappijen klanten beter informeren over hoe bepaalde beslissingen (bijvoorbeeld claims of premieberekeningen) tot stand komen. Dit bevordert niet alleen de klanttevredenheid, maar helpt ook bij het naleven van regelgeving, aangezien de verantwoording van beslissingen hiermee vereenvoudigd wordt. Bovendien stelt XAI verzekeringsbedrijven in staat om ethische overwegingen systematisch te integreren in hun AI-modellen, zodat deze niet alleen efficiënt, maar ook rechtvaardig zijn.

De integratie van XAI in de huidige datacultuur kan dus zowel uitkomsten bieden voor de klant, als voor de modeldocumentatie én het contact met de toezichthouder. In onze projecten bij verschillende opdrachtgevers zien we wel dat er nog een slag om de arm gehouden

moet worden, want er kunnen ook enkele nadelen verbonden zijn aan het toepassen van XAI. Er moet altijd een afweging gemaakt worden tussen de uitlegbaarheid en de prestaties van het model, zodat de uitkomsten uit het model de klant op een zo'n goed mogelijke manier kunnen faciliteren. De implementatie van XAI, maar ook het onderhouden van de systemen, kan daarnaast ook voor de nodige investeringen zorgen (zowel in mankracht als financieel). De efficiëntie waarmee modellen nu ontwikkeld en geïntegreerd zijn binnen een verzekeraar kan minder worden door XAI, omdat het openen van de 'black box' gepaard gaat met een uitbreiding van de modellen. Tot slot is het geen *plug-and-play* oplossing, maar zal er per model gekeken moeten worden naar de mogelijke kansen en risico's bij het gebruik van XAI.

De integratie van XAI in de verzekeringsindustrie biedt een veelbelovende oplossing voor de complexe uitdagingen waar deze sector momenteel mee geconfronteerd wordt. Door het verbeteren van de transparantie en begrijpelijkheid van AI-systemen, helpt XAI niet alleen om de vertrouwensrelatie met klanten te versterken, maar ook om de operationele en ethische uitvoerbaarheid van AI-toepassingen binnen de sector te waarborgen. ■



Limits to data quality?!

Solvency II requires data for Technical Provisions and for Internal Models to be ‘complete, accurate and appropriate’¹. DNB (2017) provides guidance “to ensure data quality”². These are ambitious requirements.

In real life, there are data quality issues. Many data quality issues have ‘solutions’ (Osborn (2024), Rao (2024) and Elahi (2022)). In this article, we present a few ‘limits to data quality’. Issues that we feel are root-causes, hard to resolve. We suggest that they deserve more attention and may well influence data quality expectations.

DNB (2017) describes the requirements for a data quality management framework in various areas. We identify root causes for data quality issues in three of these: (1) data architecture, (2) external providers and (3) objectives. We end with concluding comments.

DATA ARCHITECTURE

Two (related) root-cause issues related to ‘Single Source of Truth’ (see e.g. CRO Forum, 2020).

Issue 1: Redman (2016) writes:

The data they need has plenty of errors, and in the face of a critical deadline, many individuals simply make corrections themselves to complete the task at hand. They don’t think to reach out to the data creator, explain their requirements, and help eliminate root causes.

This is not an easy issue to resolve. The root-cause could be lack of time for the proposed feedback loop to the data creator.

Issue 2: A Reddit question: “What are the most challenging data quality issues you face frequently?”³ Wiki702 replies:

Zero schema enforcement with little to none existent governance.

Not all internal users feel fully comfortable surrendering their data models (‘schema’) to a central authority. Wang at al. (1991) identifies ‘believability’ as the first criterion for quality data. Building an effective central database takes time, to ensure that the user trusts the database.

DATA QUALITY POLICY – EXTERNAL PROVIDERS

Data quality of external data is a challenge. DNB (2017) writes:

“With external data providers, a data delivery contract has been agreed. In this contract, it described at data element level what data is to be provided, the way quality requirements are ensured and steps to be taken if the quality requirements are not met.”

Issue 3: As for ‘what data is to be provided’, the devil is in the details. Wiki identifies a host of ‘metadata’ dimensions⁴. ISO (2015) specifies ‘syntactic quality’ as a dimension of data quality. We suggest that the use of external data also requires a helpline to the ‘scientists’ driving the measurement processes. If there is an outlier, you may need to consult the scientists.

Issue 4: As to ‘quality requirement are not met’. Data (quality) requirements are dynamic, in line with changing data demand and changing data supply. How much lead-time does a supplier give when changing their data definitions? How much influence does the insurer have on the data definitions they can get from the supplier? That influence is often limited.

DATA QUALITY POLICY – OBJECTIVES

To make business decisions on data quality requires measures / scores of data quality. When are data ‘good enough’?

DNB (2017) recommends:

From the risk assessment, you may conclude required data are missing. [...] We expect that you as insurer, on the basis of this risk assessment, quantify the impact on your Solvency II reports.

The term ‘on your Solvency II reports’ is rather broad. A specific impact focus would be the Solvency Ratio. Delegated Acts (Article 222) goes one level deeper:

change or error in the outputs of the internal model, including the Solvency Capital Requirement, or in the data used in the internal



model shall be considered material where it could influence the decision-making or the judgement of the users of that information, including the supervisory authorities.

When a breach of the Solvency Capital requirement is threatened, data need to be particularly accurate. These ‘top-level criteria’ can then be cascaded down depending on the business decision.

Issue 5: It is, almost by definition, impossible to quantify (the impact of) data completeness / accuracy. We do not have the impact of missing / inaccurate data on Solvency II reports. We do not have the impact of missing / inaccurate data on decision-makers.

For such cases, article 82 of the Solvency Directive reads:
appropriate approximations, including case- by-case approaches, may be used

Yes, we can construct approximations (say scaling) to handle ‘completeness’. And, yes, we can do sensitivity analyses to assess the risk of missing / inaccurate data. But we do not know how ‘appropriate’ these approximations are, or exactly what sensitivity analyses to do. That is why we need data. The CRO Forum (2020) writes ‘Data quality is [...] a risk in itself: an operational risk’.

CONCLUDING COMMENTS

DNB (2017) writes: “*DNB expects that at all times insurers strive to attain the data quality requirements for critical data elements*”. This could be viewed as a dynamic, ‘holy grail’ interpretation of data quality requirements.

Data quality is a journey, with the following lessons learned:

- In developing new regulatory requirements, regulators and insurance companies need to cooperate to ensure that they are ‘doable’. It helps if data used for reporting can be used for steering, as interests of regulators and insurance companies are aligned.
- It is important to have a lead-time for new regulations to be implemented. This is needed to implement the new regulations, and to manage expectations about what can be delivered when. Initially, approximations may need to be used.
- If regulator / insurance company can identify data that are more ‘complete, accurate or appropriate’, the insurer is expected to use these data (‘comply or explain’).

As for the end result of the journey: we will be in data heaven, full of perfect quality data. ■

References

Atlan, 2024, “Data Quality vs Data Accuracy: 15 Key Differences To Know”, Data Quality vs Data Accuracy: 15 Key Differences To Know.

CRO Forum, 2020, “Data Quality in the Insurance Sector”.

De Nederlandsche Bank, 2017, “Guidance beheersing Solvency II data kwaliteit door verzekeraars”.

Elahi, E., May 23, 2022, 12 Most Common Data Quality Issues and Where Do They Come From, 12 Most Common Data Quality Issues – Data Ladder.

ISO (2015), “Data quality – Part 8: ‘Information and Data Quality: Concepts and Measuring’”. ISO 8000–8.

Osborn, T., Updated May 08 2024, “8 Data Quality Issues and How to Solve Them”, 8 Data Quality Issues And How To Solve Them.

Rao, S., Nov 29, 2024, 10 Common Data Quality Issues (And How to Solve Them), 10 Common Data Quality Issues (And How to Solve Them).

Redman, T., 2016, “Bad Data Costs the U.S. \$3 Trillion Per Year”, Bad Data Costs the U.S. \$3 Trillion Per Year.

Wang, Y.R., and L.M. Guarascio, 1991, “Dimensions of Data Quality: toward Data Quality by Design”, IFSRC Discussion Paper #CIS–91–06.

1 – The Directive (2009 /139/EC, Articles 82 and 121) contains data quality requirements. The Delegated Acts (2015/35) develops these in Article 222 and others (see CRO Forum, 2020).

2 – Quotes translated from Dutch.

3 – What are the most challenging data quality issues you face frequently? : r/dataengineering
https://www.reddit.com/r/dataengineering/comments/1c8dy1r/what_are_the_most_challenging_data_quality_issues/

4 – Metadata – Wikipedia
<https://en.wikipedia.org/wiki/Metadata>

T. van der Meer AAG works as Senior Actuary at Achmea. The article is written on a personal title.





Harnessing the power of AI for actuaries

Ron Richman is chief actuary at Old Mutual Insure, where he is responsible for oversight of all actuarial activities in the OMI Group. He is a Fellow of the Institute and Faculty of Actuaries (IFoA) and the Actuarial Society of South Africa (ASSA), holds practicing certificates in Short Term Insurance and Life Insurance from ASSA, and a Masters of Philosophy in Actuarial Science from the University of Cape Town. He spoke to Jennifer Baker on the subject of AI and what it means for the actuarial profession.



Some insurance CEOs refer to AI as a bubble. What is your view on this? Is this something we need to be thinking about long term?

'I think this is very similar to an exponential curve. At the beginning, it moves quite slowly, but the impacts become greater and greater with time. So there's a lot of hype currently around AI, a huge amount of investment, a lot of VC money pouring in, and that can lead one to think that maybe this is overblown. But I think we cannot underestimate what this is going to do in the longer term.

If you're using the best models out there, for example, if you are subscribed to ChatGPT and you're working with the O1-Pro model, it's almost unbelievable – you almost have a PhD-level scientific assistant in your pocket, or someone with an MBA willing to talk to you about your business. And in my experience, having worked with both of those types of people, and now having worked with the best large language models (LLM) available, we're starting to get scarily close to excellent human performance on a wide array of advanced tasks. So while I think there is maybe too much hype right now. I think we need to be very thoughtful and introspect about what the future will look like and how we make this successful for our companies, our staff, our teams and wider society.'

YOU DON'T WANT TO OVERREGULATE AND STIFLE INNOVATION. YOU DON'T WANT TO OVERREGULATE AND STIFLE INNOVATION.

How do you view the European AI Act? Do you think it's too restrictive? Does it hamper innovation? Because this is the constant refrain we sometimes hear from those who are critical of it. But the EU AI Act takes a very risk-based approach. Do you think that's the right approach?

'I think overall, the European approach is heavier on regulation. And I think if this can be implemented successfully, what the European Act will do is really ensure that you've got AI that's very well aligned with society and aligned with the sorts of goals that we expect. I think it can slow down innovation. I've seen in recent weeks some discussion of which models might be exempted. So knowing exactly where the risk lies within the various different types of models that the Act is trying to regulate is absolutely key. You don't want to overregulate and stifle innovation. I think if you look at a few of the other newsworthy items coming out of Europe in the last few weeks – for example, the significant investment in training a European LLM, or the investments into various different European AI companies – this is all positive. I think a balance needs to be found. Perhaps the European way of doing things has tended a little bit towards extra or more regulation in the past, but I think overall, the balance that we seem to be heading towards is a good thing.'

Turning specifically to the actuarial profession, what area or use purposes of AI do you think could be most transformational and most useful for actuaries?

'So I think all the hype is around the way of referencing or speaking to large language models through a chat interface. I think what we mustn't neglect is what I like to call narrow AI models, which are AI models that are specifically built for a particular purpose.

Imagine you're building a model that's excellent at actuarial pricing across a range of datasets across a range of lines of business. Those are the sorts of models I think we must focus on within the actuarial profession. You can ask a Neural Language Model (NLM) – let's say again the top tier models, like O1-Pro or the DeepSeek model – for ideas, but it will be very difficult for the model to execute a full training run against a non-life pricing dataset or set assumptions for life insurance model.

What we need to focus on are all these advances in machine learning architecture that underlie large language models. How do we take those advances and apply them to the specific niche domains where we as actuaries need to build models? Set assumptions, quantify uncertainty and do useful things. So really, what I'm quite a proponent of is steering the actuarial profession towards using narrower models that don't approach general intelligence, but rather specific actuarial intelligence. That's where we've got all of the benefits that AI can bring us, which are efficiency, scale, making much better use of our data, making more accurate predictions and quantifying uncertainty around those predictions better than we've been able to.'

I THINK THE ACTUARIAL PROFESSION NEEDS TO HAVE EXCELLENT IDEAS AROUND HOW WE AVOID OR MITIGATE THE EFFECTS OF POTENTIAL PROXY DISCRIMINATION IN THESE SORTS OF MODELS

Well, you've explained how narrow AI compares to general AI. But what sort of concerns do actuaries need to bear in mind when it comes to using AI? Where are the potential pitfalls? Here in Europe, the General Data Protection Regulation already prohibits the use of a machine to make important decisions about individual lives. So that's something I believe actuaries have already been dealing with up to now. Is there anything new in the AI Act that they might need to worry about?

'Yes. Let's look at it from the perspective first of more narrow and specific models, and then more general models.

From the perspective of more narrow and specific models, interpretability is obviously key. I think understanding the latest advances in how you can build interpretable machine learning models is really important for actuaries, so that you can understand decisions internally and be able to explain decisions made by these models externally. I don't think the discussions around the potential for proxy discrimination in machine learning or AI models is going away anytime soon. I think the actuarial profession needs to have excellent ideas around how we avoid or mitigate the effects of potential proxy discrimination in these sorts of models. And I think what we need to keep doing is pushing the limits of what these models can do for us. If we don't, that's a different sort of risk. It's a strategic risk that the work we do won't be as valuable.

I think when using large language models, everyone knows about the risks of hallucination. But I think there are more subtle risks. Even when it looks like a model isn't hallucinating, you have to spend a lot of time validating that any code written by these sorts of models is

correct. Are there subtle bugs? Has it made a subtle mistake? And I think a more general risk is, while you can get very good answers quickly out of the top tier models, that can also limit your own creativity and the limits of your own expertise as a person, I think these are risks that we need to reckon with.

How do we make sure that we're not outsourcing all of our cognitive burden off to a chatbot and our brains aren't being used to their full potential? This is something that's worrying me and a few of the colleagues that I work with. So I think understanding the universe of risks, whether it's the direct user risks resulting from using a model or the wider implications, that's where we need to spend time and effort.'

Well, that leads very clearly onto my next question, which is, do you see a role for AI in education or training of actuaries?

'I absolutely do. It's almost a paradox – if you feed the right context and the right background into a large language model it can give you fantastic suggestions on actuarial topics. Just yesterday, as a demonstration, I took O1-Pro, the top tier ChatGPT model, and asked it to design a new IBNR reserving method, and it did a pretty good job. An impressively good job, in fact. I think the key for me is, how do we get our actuaries up to a level of expertise as they start being educated, whether it's going through the university system in Europe, or the professional system like in the UK and South Africa.

How do we make sure that our next generation of actuaries is absorbing all of that information, becoming true experts, and not just outsourcing the cognitive burden to the machine? I think there'll be an unequal benefit of large language models in the future for people who've got their own expertise. I think there will be outsized benefits, because then you can really use these models to their full potential. And making sure that our next generation of young actuaries can experience those outsized benefits by being experts and of themselves, I think, is actually the core task of actuarial professions today.'

IF WE WANT TO BE SUCCESSFUL GOING INTO THE FUTURE, ACTUARIES MUST STOP BEING RECEIVERS OF THIS TECHNOLOGY, AND WE MUST START BEING CREATORS

Finally, thinking outside of the box a little bit, what are your predictions? What should actuaries bear in mind, looking to the future and considering the use of AI?

'I think that if we want to be successful going into the future, actuaries must stop being receivers of this technology, and we must start being creators. Let's take these fantastic concepts that have happened in the last, say, ten years, from the transformer model architecture, how to train transformers on huge amounts of data, how to make these models work across domains, and let's add our special actuarial touch to them, own the actuarial implications of these models, and not merely be end-users of a technology that's outside of the profession.' ■

This article has also been published in The European Actuary March 2025.



KLIMAAT

Verzekeren is vooruitzien

Niemand kan er meer omheen, ons klimaat is

veranderd, het verandert nu en het zal blijven

veranderen zolang wij broeikasgassen blijven

uitstoten en de grootschalige boskap voortzetten.

Deze klimaatveranderingen leiden wereldwijd, onder

andere, tot een toename van extreem weer. Ook in

Nederland, waar bijvoorbeeld de winterstormen en

zomerse onweersbuien, ook wel convectieve stormen

genoemd, natter worden. In de meeste gevallen

zorgen de klimaatveranderingen voor toenemende

schadelasten.



D. van Ederen PhD is student Climate-related Financial Risks en is werkzaam bij Achmea Financial Risk Management.



Dit resulteert in nieuwe uitdagingen voor de verzekerings- en herverzekeringsector, waar actuarissen al geruime tijd een cruciale rol spelen in het inschatten van klimaatgerelateerde risico's. Wellicht is een van de belangrijkste uitdagingen daarbij om de risico-managementfocus te verlengen (niet te verwarren met verleggen) naar dat wat er achter de gebruikelijke éénjaarshorizon schuilt. Een zogeheten paradigmaverschuiving. Om dergelijke verschuiving tot stand te brengen is het effectief gebleken om doelbewust reproduceerbaar bewijs te verzamelen dat het huidige paradigma doet wankelen [1]. Een taak die de wetenschappelijke methoden op het lijf geschreven is.

KLIMAATRISICO'S MANAGEN MET WETENSCHAPPELIJK ONDERZOEK

In die geest zijn er verzekeraars en andere financiële instellingen in Nederland die de banden met de academische wereld aanhalen om (de verandering in) klimaatgerelateerde risico's beter te leren begrijpen. Zo ook mijn werkgever, die zich in 2023 aansloot bij de Climate Finance Academy, een onderzoeksgroep naar klimaatgerelateerde risico's, georganiseerd vanuit het Instituut voor Milieuvraagstukken (IVM) aan de Vrije Universiteit Amsterdam. De onderzoeksgroep bestaat uit, onder andere, klimaatwetenschappers, (milieu) economen en econometristen, gedragswetenschappers en civiel ingenieurs en focust zich zowel op fundamentele als toegepaste vraagstukken gerelateerd aan klimaatrisico's voor financiële instellingen. Deze interdisciplinaire opzet sluit naadloos aan bij de vereisten om de complexe klimaatrisico's te begrijpen en te managen. Vanuit dit samenwerkingsverband hebben wij het afgelopen jaar onderzoek gedaan naar schadefuncties, ook wel vulnerability functions genoemd in de literatuur, voor winterstormschade en convectieve stormschade aan woningen.

NATUURLIJK-CATASTROFEMODELLEN

De schadefuncties geven de relatie aan tussen de intensiteit van deze stormen, en de schade aan een getroffen woning. Die intensiteit kan op verschillende manieren gemeten worden: voor de winterstormen kijkt men vaak naar de windstootsnelheid die een woning treft, bij convectieve stormen is het gebruikelijk om schade af te leiden aan de hand van hagelsteengroottes. Dit soort schadefuncties zijn onderdeel van een overkoepelend natuurlijk-catastrofemodel. Dit model kan de economische of verzekerde schade van bijvoorbeeld stormen/overstromingen/aardbevingen doorrekenen. Verzekeraars en herverzekeraars gebruiken vervolgens die resultaten om de premiehoogtes te bepalen en om de financiële buffers vast te stellen. Naast de schadefuncties bevat een natuurlijk-catastrofemodel ook een hazard- en exposuremodule. De hazardmodule beschrijft de karakteristieken van het natuurgevaar in kwestie: hebben we het over een storm of overstroming, waar en hoe vaak komt deze voor, en wat voor schade-veroorzakende condities brengt deze met zich mee? De exposuremodule bevat vervolgens alle informatie over de objecten waar men de schade over wil berekenen. In het geval van woningen gaat het dan om: de locatie van de woning, het type woning, de bouwmaterialen, bouwjaar of andere karakteristieken die de woning kwetsbaar maken voor het te bestuderen natuurgevaar. Op de intersectie tussen deze drie modules (hazard x exposure x vulnerability) bevindt zich vervolgens het risico [3].

SCHADEFUNCTIES VOOR WINTERSTORMEN EN CONVECTIEVE STORMEN

Vanuit wetenschappelijk oogpunt is het interessant om deze schadefuncties te ontwikkelen, omdat de benodigde schadedata vaak niet beschikbaar is voor academici. Verzekeraars hebben dit soort informatie wel, in de vorm van schadeclaims. Een mooi voorbeeld van hoe de industrie op haar beurt de wetenschap kan verrijken. Voor verzekeraars is het ontwikkelen van schadefuncties ook nuttig omdat daarbij een belangrijk deel van de 'black box' wordt onthuld die natuurlijk-catastrofemodellen voor hen vormen. Vrijwel iedere verzekeraar koopt namelijk deze modellen in bij derde partijen en heeft logischerwijs geen volledig inzicht in diens werking. Er zijn uiteraard legio goede redenen te bedenken waarom het inkopen van deze kennis voordelig is, zoals kostenefficiëntie. Echter bemoeilijkt het ontbreken aan volledige transparantie wel de modelvalidatie, een flexibele inzetbaarheid van het model en tot op zekere hoogte de verantwoording richting wet- en regelgevers. Desalniettemin creëert het ook een afhankelijkheid op het gebied van kwantitatieve klimaatkennis die een opbouw aan interne expertise, wat luidkeels door wet- en regelgevers wordt aangemoedigd, in de weg kan staan [4].

COMPOUND SCHADEFUNCTIES

Zelf schadefuncties ontwikkelen geeft ook de ruimte om te experimenteren met methoden die beter aansluiten bij het benodigde nieuwe paradigma. Zo weten wij bijvoorbeeld dat winterstormen en convectieve stormen onder andere natter worden, oftewel meer regen produceren [5,6]. Het is daarom van belang dat natuurlijk-catastrofemodellen dit soort ontwikkelingen kunnen verwerken om tot accurate schadeschattingen te komen. Zeker wanneer deze modellen worden ingezet om klimaatscenario's, dat zijn met behulp van klimaatmodellen gegenereerde toekomstbeelden van het klimaat en bijbehorend (extreem) weer, door te rekenen. Ondanks de geobserveerde en verwachte veranderingen in beide type stormen, blijkt dat de conventionele schadefuncties niet in staat zijn om regenvaldynamieken mee te nemen. Met onze studies, die momenteel in het peer-review proces zitten, hebben wij schadefuncties gemaakt die hier wel toe in staat zijn. Door honderdduizenden weergegelateerde schadeclaims te koppelen aan meteorologische data afkomstig vanuit het KNMI, en vervolgens verschillende combinaties van variabelen en diens transformaties te beschouwen, hebben wij een groot aantal specificaties voor schadefuncties geschat met bèta-regressies. Op basis van een model-fit-criterium hebben we vervolgens de 'beste' schadefunctiespecificatie gekozen. Het bleek dat regenvaldynamieken onmiskenbaar waardevolle informatie bevatten om schade aan een woning, ten gevolge van zowel de winterstorm als de convectieve storm, te verklaren en voorspellen [7,8]. Een recent ontstane stroming in de academische literatuur focust zich op het beschrijven van risico's ten gevolge van de samenkomst van meerdere weergegelateerde en/of klimaatgerelateerde evenementen. In het specifiek onderzoekt men combinaties waarvan de impact groter is dan de som van de individuele impacts, welke ook wel *compound weather and climate events* worden genoemd [9]. De modellen die wij ontwikkeld hebben maken daarom deel uit van een nieuwe generatie compound schadefuncties.

UNIVARIATE SCHADEFUNCTIES ONDERSCHATTEN STORMEN MET EXTREME NEERSLAG

Onze studies laten ook zien dat, wanneer een schadefunctie geen rekening kan houden met deze regenvaldynamieken, de schade aan een woonhuis wordt onderschat bij stormen met een extreme hoeveelheid neerslag. Deze onderschatting is groter voor winterstormen dan voor convectieve stormen: vanaf een 24-uurs som regenval van 50 mm (~herhalingstijd van 50 jaar [10]) gaat het over tientallen procenten voor de winterstormen. In vergelijking, voor de convectieve stormen is dit rond de 10 procent voor een 24-uurs som regenval van 100 mm (~herhalingstijd van 100 jaar [10]).

KANSEN VOOR ACTUARISSEN

Conventionele structuren, zoals wet- en regelgeving die de risico-managementfocus op een éénjaarshorizon legt, of een jarenlange afhankelijkheid van derde partijen op het gebied van kwantitatieve klimaatkennis, kunnen het lastiger maken om de risico's van de klimaatverandering ten volle te onderkennen. Desalniettemin hebben actuarissen alle competenties in huis om de noodzaak tot de paradigmaverschuiving te herkennen én deze ook te verwerkelijken. Bierviltjesberekeningen, zoals Erik Winands die in het februarinummer van De Actuaris aandroeg, kunnen een interessant middel zijn om deze paradigmaverschuiving op gang te brengen. Temeer omdat deze snel tot een orde van grootte van het te besturen fenomeen kunnen te komen. En het is juist die snelheid die erg hard nodig is in een tijd waarin 'the window of opportunity' om de klimaatverandering tegen te gaan, snel sluit. ■

Referenties

[1] IPCC, 2021: *Climate Change 2021: The Physical Science Basis. Contribution of Working Group I to the Sixth Assessment Report of the Intergovernmental Panel on Climate Change* [Masson-Delmotte, V., P. Zhai, A. Pirani, S.L. Connors, C. Péan, S. Berger, N. Caud, Y. Chen, L. Goldfarb, M.I. Gomis, M. Huang, K. Leitzell, E. Lonnoy, J.B.R. Matthews, T.K. Maycock, T. Waterfield, O. Yelekçi, R. Yu, and B. Zhou (eds.)]. Cambridge University Press, Cambridge, United Kingdom and New York, NY, USA, In press, doi: <https://www.cambridge.org/core/books/climate-change-2021-the-physical-science-basis/415F29233B8BD19FB55F65E3DC67272B>.

[2] Taylor, E. W., & Cranton, P. (2012). *The handbook of transformative learning: Theory, research, and practice*. John Wiley & Sons.

[3] Kron, W. (2005). Flood risk= hazard• values• vulnerability. *Water international*, 30(1), 58–68.

[4] DNB (2023). Gids voor de beheersing van klimaat- en milieurisico's

[5] KNMI, 2023: KNMI'23klimaatscenario's voor Nederland, KNMI, De Bilt, KNMI-Publicatie 23–03.

[6] KNMI 2025: De staat van ons klimaat 2024; Nederlands weer in tijden van klimaatverandering, KNMI, De Bilt, KNMI-Publicatie 25–01

[7] van Ederen et al. (2025). A high-resolution compound vulnerability function for European winter storm losses (Preprint). <https://scity.org/articles/activity/10.21203/rs.3.rs-5618142/v1> (Manuscript submitted for publication)

[8] van Ederen et al. (2025). A high-resolution compound vulnerability function for severe convective storm storm losses. (Manuscript submitted for publication)

[9] Zscheischler, J., Martius, O., Westra, S., Bevacqua, E., Raymond, C., Horton, R. M., ... & Vignotto, E. (2020). A typology of compound weather and climate events. *Nature reviews earth & environment*, 1(7), 333–347.

[10] STOWA. (2019). *Neerslagstatistiek en-reeksen voor het waterbeheer 2019* (STOWA Rapport 2019–19). STOWA. <https://www.stowa.nl/sites/default/files/assets/PUBLICATIES/Publicaties%202019/STOWA%202019-19%20neerslagstatistieken.pdf>



Multi-stage stochastische modellering adresseert tekortkomingen van Markowitz' portfolio-optimalisatie

Het construeren van optimale portfolio's is een fundamenteel vraagstuk in de financiële wereld. Portfolio-optimalisatie richt zich op het bepalen van een kapitaalverdeling over financiële activa die de risico-rendementverhouding optimaliseert. Hiermee wordt gestreefd naar een risico-efficiënt portfolio: een verdeling waarbij, gegeven een bepaald risiconiveau, geen enkele andere combinatie van activa een hoger verwacht rendement oplevert. Een grote uitdaging bij bestaande methoden voor portfolio-optimalisatie is het effectief omgaan met de onvoorspelbaarheid en dynamiek van financiële markten. Deze factoren maken modellering complex en vereisen inventieve technieken. Om deze uitdaging te adresseren, heb ik in mijn scriptie een multi-stage stochastisch model voor portfolio-optimalisatie ontwikkeld. In dit artikel bespreek ik de belangrijkste bevindingen van mijn onderzoek.

L. Beute MSc is werkzaam als trainee bij TVM Verzekeringen en winnaar van de Johan de Witt prijs 2024.



HISTORISCHE CONTEXT

In 1952 introduceerde Harry Markowitz het mean-variance (MV) model; de basis voor het zogeheten portfolio selectieprobleem (PSP). In het PSP wordt getracht een portfolio te bepalen die risico-efficiënt is. Het MV-model probeert dit te bereiken door portfolio's te selecteren waarbij de verhoudingen tussen (co)variantie en verwacht rendement optimaal zijn.

Hoewel Markowitz' theorie baanbrekend was, is deze in de loop der tijd ook bekritiseerd. Een belangrijke kritiek is dat de versimpelde benadering in het MV-model de complexiteit van de financiële markten niet goed weergeeft. Hierdoor zou het model niet goed in staat zijn om praktische omstandigheden effectief te adresseren. Een voorbeeld hiervan is de aanname dat rendementen normaal verdeeld zijn; een veronderstelling die inmiddels meermaals is weerlegd en niet langer als valide wordt beschouwd. Bovendien beperkt het MV-model de integratie van praktische beperkingen, zoals minimale positiegroottes en limieten op het aantal unieke activa in een portfolio. Deze tekortkomingen kunnen leiden tot significante schattingsfouten in de rendementen van de portfolio's die met het MV-model geconstrueerd worden.

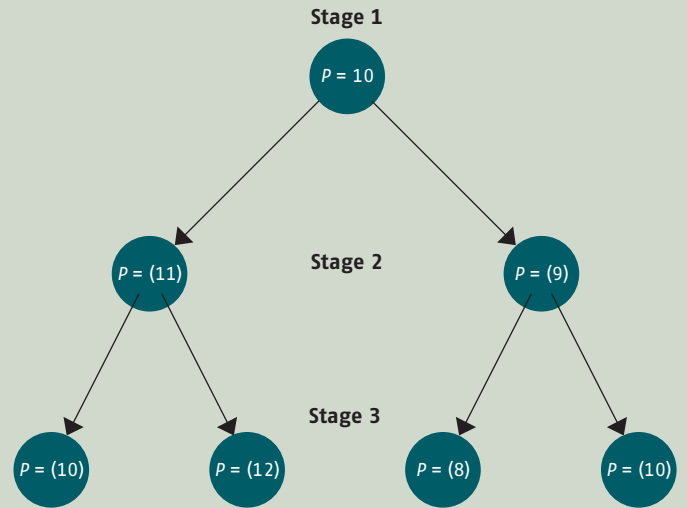
Vanwege deze beperkingen kan gesteld worden dat het MV-model suboptimaal is voor portfolio-optimalisatie. Moderne benaderingen, zoals robuuste optimalisatie en machine learning, bieden verfijningen. Een veelbelovende, maar nog relatief onbenutte methode voor het PSP is multi-stage stochastische modellering.

MULTI-STAGE STOCHASTISCHE MODELLERING

Multi-stage stochastische modellering is een methode waarbij besluitvorming plaatsvindt over meerdere tijdstappen, ook wel stages genoemd. In elke stage zijn er verschillende scenario's, ieder met mogelijke realisaties van een stochastisch proces. Het generieke doel van een multi-stage stochastisch model is het minimaliseren van de som van de initiële kosten en de verwachte toekomstige kosten die voortkomen uit de genomen besluiten, afhankelijk van de mogelijke realisaties van de stochast.

Onderliggend aan deze modellen ligt vaak een zogeheten scenario tree, die de verschillende stages en scenario's representeert. In mijn modellering van het PSP ga ik ervan uit dat de toekomstige prijs van financiële instrumenten de stochast is waarop de beslissingen worden gebaseerd. De afbeelding hierna biedt een visuele weergave van een kleine scenario tree in de context van het PSP. Hierin stelt P de stochast voor die de prijs van een arbitrair financieel instrument beschrijft. Stage 1 kan worden geïnterpreteerd als het huidige moment, waarbij de prijs bekend is. Stage 2 en stage 3 representeren momenten in de

toekomst. Tussen de haakjes wordt de mogelijke realisatie van de stochast P weergegeven in een specifiek prijsscenario van een bepaalde stage.



Aan de hand van deze scenario trees heb ik een model ontwikkeld dat een optimaal portfolio construeert. Dit portfolio minimaliseert de som van de risicomaatstaf Conditional Value at Risk (CVaR) in de eerste stage en de verwachte CVaR in de toekomstige stages. Tegelijkertijd wordt dit portfolio geselecteerd zodat een minimaal vereist verwacht rendement wordt behaald. Het resulterende model heeft de structuur van een multi-stage stochastisch model, waarbij in iedere stage de mogelijkheid bestaat om het portfolio te herbalanceren op basis van de gerealiseerde prijsscenario's. In mijn scriptie laat ik zien dat deze structuur de flexibiliteit biedt om praktische factoren toe te voegen aan de modellering, zoals het vereisen van minimale positiegroottes en limieten op het aantal unieke activa. Daarnaast biedt het model ruimte om de toekomstige prijzen en daarmee rendementen flexibel te modelleren. Op deze manier worden de eerdergenoemde tekortkomingen van het MV-model effectief geadresseerd.

COPULA'S: VERBETERING IN MODELLERING VAN RENDEMENTEN

Bij multi-stage stochastische modellering is de nauwkeurigheid van de modellering van de stochast van cruciaal belang. Binnen de context van het PSP vormt dit een uitdaging, aangezien financiële rendementen een complexe datastructuur kennen, waarbij ze – zoals eerder besproken – niet normaal verdeeld zijn en marginaal variëren. Voor deze structuur dient een multivariate kansverdeling te worden geschat.

Om deze uitdaging te overwinnen, heb ik een parametrische copula-methodiek ontwikkeld. In mijn scriptie toon ik aan dat deze methodiek in staat is om prijsscenario's van hoge kwaliteit te genereren. Copula's zijn verdelingsfuncties die de afhankelijkheden tussen meerdere variabelen kunnen modelleren, terwijl de marginale kansverdelingen van de variabelen afzonderlijk blijven. Dit maakt copula's bijzonder geschikt voor het modelleren van verdelingen van financiële rendementen, omdat ze de correlaties tussen variabelen vastleggen zonder de specifieke vorm van de marginaal verdeelde variabelen te beïnvloeden.

RESULTATEN

Om de prestaties van het multi-stage stochastische model te testen, heb ik een backtest-experiment uitgevoerd op basis van de Dow Jones Index (DJI); een van de belangrijkste aandelenindices in de Verenigde Staten. Met het multi-stage stochastische model heb ik portfolio's samengesteld uit DJI-assets en de rendementen vergeleken met de DJI zelf. De prijsscenario's die in het model zijn gebruikt voor deze assets, zijn gegenereerd volgens de parametrische copula-methodiek. De

periode die is gebruikt voor de schatting is 2010–2018; de prestaties zijn vergeleken met de resultaten van 2019. Om de robuustheid van het model te testen, heb ik gebruikgemaakt van verschillende minimale wekelijkse rendementseisen, variërend van 0,1% tot 0,4% per week—een bandbreedte die goed aansluit bij de historische wekelijkse rendementen van de DJI.

De tabel toont de resultaten van dit backtest-experiment. Een belangrijke observatie is dat, hoewel de gemiddelde gerealiseerde wekelijkse rendementen van alle portfolio's dicht bij elkaar liggen, de risicomaatstaven standaarddeviatie en CVaR aanzienlijk lager zijn voor de modelgebaseerde portfolio's.

	Multi-stage stochastische model portfolio's				DJI
Minimale wekelijkse rendementeis (input)	0,10%	0,20%	0,30%	0,40%	–
Gemiddeld wekelijk rendement	0,38%	0,39%	0,42%	0,44%	0,40%
Dagelijkse 95% CVaR	–1,48%	–1,47%	–1,53%	–1,83%	–1,97%
Dagelijkse standaard deviatie	0,65%	0,65%	0,66%	0,77%	0,79%

Deze observatie impliceert dat het model in staat is om portfolio's te genereren met een verbeterde risico-efficiëntie. In de praktijk kan een verbeterde risico-efficiëntie belangrijke implicaties hebben. Neem bijvoorbeeld een pensioenfonds: door de risico-efficiëntie van de beleggingsportfolio's verder te verbeteren, kan het risico op grote verliezen worden verkleind, terwijl de stabiliteit van de pensioenuitkeringen behouden blijft. Desgelijks kunnen banken hun kapitaal efficiënter inzetten en tegelijkertijd voldoen aan wet- en regelgeving

HOE NU VERDER?

Een belangrijk aspect dat in mijn scriptie uitgebreid wordt besproken, maar in dit artikel minder aan bod komt, is de complexiteit van multi-stage stochastische modellering. Dergelijke modellen groeien snel in omvang, wat ze computationeel intensief maakt. Bepaalde modelleringskeuzes kunnen bovendien leiden tot het gebruik van geheel tallige variabelen, zoals het geval is bij het multi-stage stochastische PSP-model dat ik in mijn scriptie ontwikkel. Dit verhoogt de computationele complexiteit verder.

Door de unieke structuur van dit type model binnen de context van het PSP, bestaan er momenteel geen standaard efficiënte oplossingsmethoden voor het verkrijgen van exacte oplossingen. Dit vereist een beperking van het aantal scenario's om tot oplossingen te komen. In mijn scriptie introduceer ik een algoritme dat deze bottleneck helpt te doorbreken, maar dit kan verder verfijnd worden. Tegelijkertijd kijk ik vooruit: met de toenemende computationele capaciteiten en nieuwe algoritmische ontwikkelingen kunnen in de toekomst verdere optimalisaties mogelijk zijn. Dit zou het model in staat stellen meer scenario's te verwerken, wat naar verwachting prestaties verder zal verbeteren. ■



PENSIOEN

Dynamisch voorsorteren op Wtp-renteafdekking

Hier volgt een filemelding. Bij de aanstaande transitie naar de nieuwe pensioenregeling wordt een opstopping voorzien op de financiële markten. Aan het begin van 2026/27 sluiten veel pensioenfondsen aan in de rij om hun renteafdekking in lijn te brengen met de nieuwe pensioenregeling. U wilt de file vermijden. Maar hoe doet u dat?

Veel pensioenfondsen zijn voornemens om in het begin van 2026 of 2027 over te stappen op de nieuwe pensioenregeling (de solidaire of flexibele pensioenregeling) en bestaande pensioenopbouw in te varen. Dit heeft impact op de renteafdekking. In de huidige regeling heeft de renteafdekking een collectief karakter. In de nieuwe regeling daarentegen sluit de renteafdekking aan bij de risicohouding van de verschillende leeftijdscohorten. Hierdoor kent de nieuwe regeling een relatief hoge renteafdekking voor ouderen (laag risicoprofiel) en een relatief lage renteafdekking voor jongeren (hoog risicoprofiel). Dit betekent dat in de nieuwe regeling veel kortlopende en weinig langlopende renteafdekking benodigd is.

Het is denkbaar dat de overstap op de renteafdekking in de nieuwe pensioenregeling veel tijd kost:

- Afhankelijk van de fonds-specifieke situatie kunnen de verschillen tussen de renteafdekking in de huidige versus nieuwe regeling groot zijn. In dat geval dienen grote volumes te worden verhandeld om te overstap te maken.
- Veel pensioenfondsen maken gebruik van drukke transitiemomenten (1 januari 2026 en 1 januari 2027) en doen daardoor gelijktijdig een beroep op liquiditeit die met name voor lange looptijden beperkt beschikbaar is.

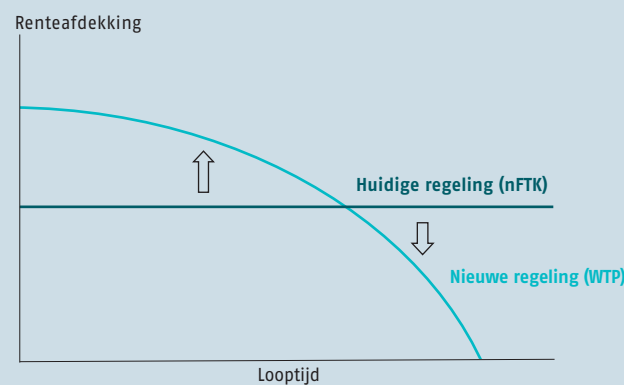
Een nadelig gevolg hiervan is dat bij aanvang van de regeling, de feitelijke afdekking van het renterisico niet aansluit bij de in de nieuwe regeling benodigde afdekking. Hierdoor worden bepaalde groepen tijdelijk te veel of te weinig blootgesteld aan renterisico.

Drs. L. Nagel CFA CAIA (links) is werkzaam als ALM-strateeg bij pensioenfondsen Rail & OV.

Drs. P.J.T. Marres AAG is werkzaam als data actuaaris bij QuantSense B.V.



Pensioenfondsen kunnen voorsorteren door reeds vóór transitiedatum de renteafdekking geheel of gedeeltelijk in lijn te brengen met de nieuwe pensioenregeling. Hiermee bedoelen we dat in de huidige regeling de renteafdekking per looptijd alvast in lijn gebracht wordt met de benodigde afdekking per looptijd in de nieuwe regeling, zie figuur 1. Zo kan worden voorkomen dat de renteafdekking vlak na invaren, tegelijkertijd met andere fondsen, dient te worden bijgesteld.



Figuur 1: Voorbeeld van de mogelijke gevolgen van de overstap naar de nieuwe regeling op de renteafdekking.

BENODIGDE RENTEAFDEKKING IN DE WTP

Om te kunnen voorsorteren dienen de beoogde renteafdekkingspercentages per looptijd in de nieuwe regeling bekend te zijn. Voor fondsen die al zijn ingevaren, volgen deze percentages uit de (pensioen)administratie. Pensioenfondsen die nog niet zijn ingevaren en nog niet beschikken over een administratie, kunnen de afdekkingspercentages inschatten. Deze inschatting is met name afhankelijk van de volgende ingrediënten:

- I. De uitgangspunten voor invaren. Hiermee bedoelen we de verdeling van het fondsvermogen over collectieve reserves en individuele deelnemers (en daarmee ook over leeftijdscohorten).
- II. De benodigde renteafdekking per leeftijdscohort op basis van het beleggingsbeleid in de nieuwe pensioenregeling.
- III. Een aanname voor de invaardekkinggraad. De invaardekkinggraad is van belang, omdat pensioenfondsen in de nieuwe regeling beleggingen afdekken in plaats van verplichtingen. Hierdoor dekken zij na invaren ook een eventueel dekkingsgraadoverschot af. Het ligt voor de hand om de huidige dekkingsgraad als aanname voor de toekomstige invaardekkinggraad te gebruiken.

VOORSORTEREN OP WTP-RENTAFAFDEKKING LEIDT TOT DYNAMISCH BELEID

De inschatting van de in de nieuwe regeling benodigde renteafdekking (Wtp-renteafdekking) fluctueert over tijd. Deze inschatting is immers afhankelijk van de dekkingsgraad (aanname voor de invaardekking-

graad), die op zijn beurt afhankelijk is van ontwikkelingen op financiële markten. Het is daarom van belang om periodiek te monitoren of de feitelijke renteafdekking nog steeds in voldoende mate aansluit bij de Wtp-renteafdekking, bijvoorbeeld met behulp van dekkingsgraadgrenzen. Het ligt voor de hand om de Wtp-afdekking te herrekenen als de dekkingsgraad zich buiten een bepaalde bandbreedte begeeft, omdat in dat geval de Wtp-afdekking waarschijnlijk gedaald of gestegen is. Naast dekkingsgraadveranderingen kunnen ook andere aspecten aanleiding geven tot herrekening van de Wtp-afdekking, bijvoorbeeld wijzigingen van de uitgangspunten voor invaren of ontwikkelingen in het deelnemersbestand.

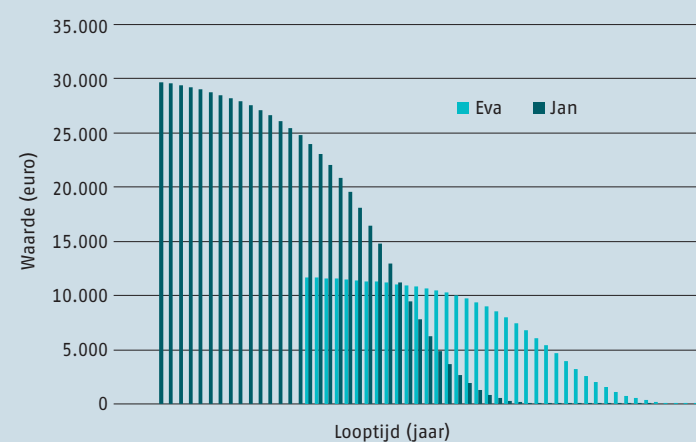
REKENVOORBEELD

Laten we een pensioenfonds 'Vooruitziende blik' (hierna: 'pf VB') bekijken met twee deelnemers: Eva (50 jaar) en Jan (65 jaar). Pf VB heeft nu nog het (collectieve) renterisico voor 60% afgedekt. De voorzieningen van Eva en Jan worden bij een dekkingsgraad van 130% door pf VB ingevaren volgens de standaardmethode¹, zie tabel 1.

	Eva	Jan
Voorziening (nFTK)	259	507
Invaarbonus	84	127
Persoonlijk Pensioen Vermogen (Wtp)	343	634

Tabel 1: Eva en Jan varen in bij een dekkingsgraad van 130% (bedragen in EUR x 1.000,-)

De in de tabel weergegeven Persoonlijke Pensioen Vermogens van Eva en Jan worden op basis van actuariële factoren omgezet naar een kasstroom voor zowel ouderdomspensioen als partnerpensioen. In de nieuwe regeling is de afdekking van deze kasstromen leeftijdsafhankelijk: naarmate een deelnemer dichterbij zijn/haar pensioendatum komt, wordt een hoger percentage van de kasstromen afgedekt. Bij pf VB betekent dit dat in de nieuwe regeling voor Eva 40% van haar kasstroom wordt afgedekt en voor Jan 70%, zie figuur 2.



Figuur 2: af te dekken kasstromen

De figuur laat zien dat de korte looptijden worden gedomineerd door oudere deelnemers met hoge afdekkingspercentages, in het geval van pf VB zijn dat de kasstromen van Jan. In tabel 2 zijn de af te dekken kasstromen van Eva en Jan bij elkaar opgeteld en uitgedrukt als percentage van de kasstromen in de huidige regeling. Hieruit volgt dat om voor te sorteren, de afdekking op korte looptijden dient te worden verhoogd en op lange looptijden verlaagd.

looptijd	Huidige regeling (nFTK)	Nieuwe regeling (WTP)
5 – 10 jaar	60%	80%
11 – 20 jaar	60%	86%
20 – 30 jaar	60%	69%
30 – 40 jaar	60%	57%
40 – 50 jaar	60%	50%
50 – 60 jaar	60%	48%

Tabel 2: Effect van voorsorteren op Wtp-renteafdekking

VAN REKENSTRAAT NAAR REKENSNELWEG

Om per leeftijdscohort de nieuwe renteafdekking onder de nieuwe regeling (Wtp) te berekenen, doorlopen we achtereenvolgens de volgende stappen:

1. opschonen, koppelen en prepareren van de data;
2. omzetten van de opgebouwde pensioenaanspraken –en rechten naar persoonlijke pensioenvermogens ('invaren') volgens de standaardregel;
3. omzetten van de persoonlijke pensioenvermogens van alle deelnemers naar af te dekken Wtp-kasstromen;
4. alle individuele kasstromen uit punt 3 aggregeren we tot één collectieve af te dekken Wtp-kasstroom. Deze collectieve Wtp-kasstroom zetten we vervolgens af tegen de collectieve nFTK-kasstroom om te komen tot een renteafdekking per leeftijdscohort.

Spreadsheetsoftware als Microsoft Excel blijkt ontoereikend om deze rekenstraat in afzienbare tijd te doorlopen. De bottleneck zit met name in stap 3. Door de data te normaliseren en kasstromen als vectoren te verwerken in een efficiënte programmeeromgeving (Python) kan stap 3 binnen vijf minuten worden doorlopen met een deelnemersbestand met diverse aanspraken voor elk van de 100k+ deelnemers.

CONCLUSIE

De wijze waarop het renterisico wordt afgedekt in het nieuwe stelsel (Wtp) is fundamenteel anders dan we gewend waren. Onder Wtp wordt de renteafdekking leeftijdsafhankelijk. In de praktijk betekent dit dat bij oudere deelnemers een groter deel van het vermogen wordt beschermd tegen renteveranderingen dan bij jongere deelnemers. In dit artikel hebben we laten zien hoe een pensioenfonds de drukte na populaire invaarmomenten kan voorkomen door reeds voor invaren de renteafdekking in lijn te brengen met de nieuwe pensioenregeling. ■

1 – Rekening houdend met het initieel vullen van het minimum vereist eigen vermogen en de operationele reserve.

per 1 februari

NIEUWE LEDEN

M.W.N. Zeeman MSc AAG (Max)	Lid AAG
S.A.A. Dechesne AAG (Suzanne)	Lid AAG
J.F. ter Braake AAG (Jelte)	Lid AAG
B. Belkadi MSc AAG (Bilal)	Lid AAG
A. Zariouh MSc AAG (Abdel)	Lid AAG
M. Wilke AAG (Max)	Lid AAG
Y. Li AAG (Yanan)	Lid AAG
C. Swanepoel (Carien)	Lid geaffilieerd
K. de Boer (Klaartje)	Lid geaffilieerd
D. Sikma MSc (Dennis)	Lid student
E. de Heer (Emma)	Lid student
L.J. Gerrits MSc (Lotte)	Lid student
Z.X.S. Paraguas MSc (Zeus)	Lid student
S. Johannes MSc (Swen)	Lid student

per 1 maart

E.A.M. Simons MSc AAG (Esmée)	Lid AAG
M.S. Schaeckelhoff AAG (IA/BE) (Maren)	Lid AAG
S.A.M. Plat MSc AAG (Sharon)	Lid AAG
J. Hablous MSc AAG (Jos)	Lid AAG
M. Peetoom MSc AAG (Mick)	Lid AAG

AGenda

Lessons learned: Dubbele materialiteit en Waardeketen bij verzekeraars | 8 mei | Verbond van Verzekeraars (Den Haag)

Kring Young Actuaries AG
16 mei | Johan de Witt huis (Utrecht)

Algemene Ledenvergadering
3 juni | Spant! (Bussum)

AG Jaarcongres 2025
Barometer van de pensioentransitie
– hoe staan we ervoor? | 3 juni | Spant!
(Bussum)

Kijk voor meer informatie over de bijeenkomsten van het AG in de online agenda:

**EMAS-DIPLOMA'S UITGEREIKT**

Vrijdag 14 maart 2025 ontvingen 26 studenten hun diploma Actuaris. Tijdens de feestelijke bijeenkomst in Kasteel Woerden reikte Karima Attrach, opleidingscoördinator EMAS, de diploma's uit. Hartelijk gefeliciteerd allemaal!

De geslaagden

Matthias Borst, Jelte ter Braake, Nick de Bruijn, Thijn Burgerhof, Ekaterina Dimitrova, Arjan Dwarshuis, Robbert-Jan Groen, Joost Groeneveld, Jos Hablous, Martina Kamburova, Ariadne Koletsos, Aleksandra Lalatovic, Cas Lutz, Feddrick van der Meer, Wouter Mertens, Fabian Polman, Erick Ferreira Reyes, Geert Roekaerts, Jelle Schakenraad, Annelore Schilstra, Lilian Shahnazari, Wouter de Voogd, Lisanne de Vos, Max Wilke, Kim Wittekoek en Abdel Zariouh.



Op de foto (vlnr): Wilma Venmans (voorzitter Centrale Examencommissie), Edwin Roebersen (bestuurslid), Loudina Erasmus (bestuurslid), Koos Gubbels (academic director EMAS), Kim Wittekoek, Aleksandra Lalatovic, Robbert-Jan Groen, Ariadne Koletsos, Ekaterina Dimitrova, Erick Ferreira Reyes, Martina Kamburova, Jelte ter Braake, Thijn Burgerhof, Joost Groeneveld, Matthias Borst, Cas Lutz, Feddrick van der Meer, Arjan Dwarshuis, Abdel Zariouh, Wouter Mertens, Jos Hablous, Fabian Polman, Annelore Schilstra, Wouter de Voogd, Jelle Schakenraad, Max Wilke, Geert Roekaerts en Lisanne de Vos.

Bijdragen aan de komende thema's van De Actuaris?

Beste lezer,

Hierbij presenteren wij de thema's voor de komende nummers. Mocht je een bijdrage overwegen, of bepaalde suggesties of wensen hebben, dan horen wij deze graag! Aarzel dus niet om contact op te nemen met de redactie. Wij zijn erg benieuwd naar je reactie!

Juni 2025: Pensioen

Pensioen is al tijden een hot topic. Overal hoor of lees je iets over het nieuwe pensioenstelsel. Iedereen vindt er wat van of roept er iets over. Inmiddels bereiden pensioenfondsen zich voor op de invoering. De transitieplannen worden per 1 januari 2025 opgeleverd en voor 1 juli 2025 staan de implementatieplannen op de planning. We zitten er met de neus bovenop. In deze editie gaan we het hier uitgebreid over hebben.

Thematrekkers: Sanne van Helvert en Pieter Bouwknecht

Jaargang 2025-2026

Voor de jaargang 2025 – 2026 is de redactie volop bezig met de voorbereiding van de thema's. In het juninummer worden die vermeld.

De redactie:

Pieter Bouwknecht (pieter.bouwknecht@nn.nl)
Robin Cats (robincats@gmail.com)
Salima El Khababi (salima.elkhababi@milliman.com)
Rens Garssen (rens.garssen@oliverwyman.com)
Koos Gubbels (Koos.Gubbels@achmea.nl)
Lars Janssen (lars.janssen@pwc.com)
Anne Joosten (anne.joosten@nl.ey.com)
Elke Op het Veld (elke.op.het.veld@sprenkels.nl)
Sanne Schelfhout-van Helvert (sanne.van.helvert@aaa-riskfinance.nl)



colofon de actuaris – jaargang 32 – nr 4 – magazine van het Koninklijk Actuariel Genootschap – ISSN 0929-4562

redactie

Pieter Bouwknecht
Robin Cats
Salima El Khababi
Rens Garssen
Koos Gubbels
Lars Janssen
Anne Joosten
Elke Op het Veld, hoofdredacteur
Sanne Schelfhout-van Helvert
Frank Thooft

eindredactie

Frank Thooft

gemandateerd uitgever

Maarten van Meerten

contact

Koninklijk Actuariel Genootschap
Groenewoudsedijk 80
Pascale Mandjes-Heese
3528 BK Utrecht
E redactie@actuarielgenootschap.nl
T 030 – 686 61 50

vormgeving

Stahl Ontwerp

druk

Print Power Media

kopij

Voor het volgende nummer (juni 2025) dient de kopij uiterlijk **6 mei 2025** digitaal ingeleverd te worden bij de redactie: redactie@actuarielgenootschap.nl.

Auteursinstructies staan op <https://www.actuarielgenootschap.nl/over-het-koninklijk-actuariel-genootschap/magazine-de-actuaris>

De redactie behoudt zich het recht voor artikelen te weigeren.

achtergrond

De Actuaris verschijnt vijf keer per verenigingsjaar met interviews, nieuws, informatie en opinievormende artikelen die van belang kunnen zijn voor de actuariële beroepsgroep en degenen die door opleiding en of interesse het actuaariaat na staan. Het overnemen en vermenigvuldigen van artikelen met bronvermelding is toegestaan na toestemming van de redactie. Alle artikelen uit deze uitgave worden online beschikbaar gesteld in de Kennisbank op de website van het AG.

disclaimer

Hoewel aan de totstandkoming van 'De Actuaris' de uiterste zorg is besteed, aanvaarden de auteur(s), redacteur(en) (Redactie) en het bestuur AG, alsmede de uitgever(s), geen enkele aansprakelijkheid voor eventuele fouten en of onvolkomenheden, noch voor de gevolgen daarvan.

'De Actuaris' wordt uitgegeven in opdracht van het bestuur AG. De in het tijdschrift voorkomende meningsuitingen mogen echter niet worden gezien als de officiële zienswijzen van de Redactiecommissie en/of het bestuur AG, tenzij zulks uitdrukkelijk is vermeld.





Kijk voor opleidingen op
www.actuarieelinstituut.nl