



ISSN 0803-8953

Tilfeldig gang

Nr. 1, årgang 37

Mai 2020

Utgitt av Norsk statistisk forening

Redaksjonelt

Fra NFR-lederen	1
Fra redaksjonen	2

Tankenøtter

Pusleri nr. 53	3
Pusleri nr. 52 (løsning)	4

Møter og konferanser

Oppstart NSF Stavanger	7
Medlemsmøte NSF Trondheim	8

Artikler

Monitoring the Level and the Slope of the Corona	9
På leting etter sannheten bak kvantemekanikken	22
Tilbakeblikk: Den første nordiske konferanse i matema- tisk statistikk, Aarhus 1965	28

Kunngjøringer

Kontingent	31
----------------------	----

Meldinger

Nytt fra NTNU	32
Nytt fra UiB	32
Nytt fra Matematisk institutt ved Universitetet i Oslo .	32

Fra NFR-lederen

Hei alle NFR-medlemmer,

Alt vel med dere og nærmeste familie/venner i disse corona-tider håper jeg – alt vel her.

Vi lever alle uvanlige liv for tiden. Mennesket er jo et sosialt vesen – og sosiale medier ivaretar bare en flik av det. Jeg har i lengre tid hatt barn som studenter i utlandet – og når vi samles om sommeren på hytta har vi et begrep 'beyond Skype' – det betyr en god bamse-klem. Det er behov for gode klemmer – og den tid kommer tilbake.

Vi føler vel alle savnet av familie-medlemmer, jobb-kolleger og venner – samt for egen del, savnet av en Kreta-tur i mai. Jeg har, som gammel 'low-tech' professor, blitt påtvunget digitalt basert undervisning – og føler meg ikke trygg på at studentene er begeistret. Personlig ser jeg frem til normale undervisnings-tilstander.

Statistikken i coronaens tid – lever faktisk i beste velgående. 'Usikkerhet' er vel det mest brukte ordet i samfunns-debatten siste måned. Statistiske modeller i epidemiologi er på alles lepper – og diskusjonen går høyt om prøvetaking og test-prosedyrer. Det er faktisk mer spennende å være statistiker enn noen gang. Jeg regner med at vi alle har lest 'antall smittede', 'antall innlagt' og 'antall på intensiv' med stor interesse siste måned – og konstruert vår statistiske modeller enten mentalt eller på et kladdeark.

Sommerens statistikk-høydepunkt for mange skulle være Nordstat2020 i Tromsø 23.-25. juni – og det er vel ingen overraskelse at denne konferansen, i skrivende stund, henger i en tynn tråd. Formelt er ikke Nordstat2020 avlyst/utsatt ennå, men i lesende stund er den nok det. La meg først takke ArrKom i Tromsø for glimrende arbeid så langt – de har fattet uvanlig mange beslutninger under uvanlig stor usikkerhet – som de jo som statistikere er velegnet til. Mye tyder på at vi andre får en Tromsø-tur i løpet av de nærmeste år.

Det blir en rolig sommer i lokal-miljøet for de fleste av oss. Vi har sikkert ikke vondt av å roe litt ned. Tid på Hytta, i Vogna eller i Telt i fedrelandet er fine greier – i hvertfall hvis sommer-været viser seg fra sin beste side.

God Sommer – ta vare på dere selv og andre ! HILSEN
HENNING O

Fra redaksjonen

Stein Andreas Bethuelen, Jan Terje Kvaløy og Tore Selland Kleppe | Stavanger

Som også nevnt av Henning O er dette nummeret stiftet sammen i en tid der verden er hardt rammet av SARS-CoV-2. I tillegg til å ha fått arbeids- og undervisningssituasjon fullstendig satt på hodet, er mange statistikere involvert i modelleringsarbeid relatert til viruset. Likevel har vi fått inn en god mengde bidrag til dette nummeret av TG. Tusen takk til alle som har bidratt, og spesielt takk til Jostein Lillestøl, Nils Lid Hjort, Emil Aas Stoltenberg og Inge Helland!

Jostein Lillestøl har i korrespondanse minnet oss på at TG spiller en viktig rolle i dokumentering av aktivitet knyttet til statistikkfaget i Norge og NSF's aktiviteter for ettertiden. Dette bør vi alle ha i bakhodet både når det gjelder dokumentering av tidsnære hendelser og tanker, og også de som ligger lengre tilbake i tid.

Også i tiden fremover har vi ambisjoner om å gi ut bladet ihvertfall to ganger i året (høst/vår). Sitter du på materiale av en viss lengde og som helst bør ut i løpet av kort tid, vil vi selvsagt gi ut spesialnummer utenfor den vanlige planen. Vi håper også å gjenoppta formatet der nydisputerte PhDer beskriver sin forskning. Veiledere oppfordres til å minne sine studenter på dette (det kom ingen bidrag etter siste oppfordring). Ellers er som vanlig alt mulig relevant materiale av interesse, og kan sendes til tore.kleppe@uis.no.

Pusleri nr. 53

Jostein Lillestøl | Bergen

Fem pensjonistvenner har avtalt å møtes på eldresentret en gang i uken på et bestemt tidspunkt. De pleier å ankomme hver for seg i tilfeldig rekkefølge. Tre av dem har imidlertid uvanen at dersom ingen andre er der ved ankomst, så går denne sin vei uten å komme tilbake den dagen. Hva er sannsynligheten for at de blir mange nok til å spille bridge? Hva er forventet antall fremmøtte? En av de upålitelige er strammet opp, og lover heretter å vente på de andre. Hva blir nå sannsynligheten og forventningen?

Finn uttrykket for fordelingen til antall fremmøtte X og dens forventning generelt for tilfellet med n personer, hvorav m er pålitelige og $n - m$ er upålitelige, og går sin vei dersom ingen andre er kommet før.

Pusleri nr. 52 (løsning)

Jostein Lillestøl

Oppgaven var: En kinoforestilling skal snart begynne, og åtte personer er igjen i billettøen. Billettprisen er 100 kroner per person. Billettøren tar kun kontanter, og har akkurat sluppet opp for vekslepenger. Anta at fire av de åtte har 100 kroner, mens de andre fire bare har en 200 kroner seddel, og at kundene står i tilfeldig rekkefølge.

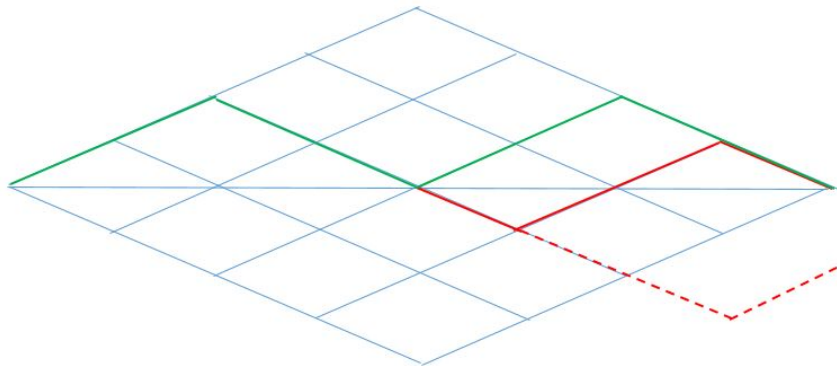
Hva er sannsynligheten for at det oppstår veksleproblem?

Hva er den generelle løsning når det er $m=2n$ kunder i køen?

Hva om det er to flere med 200 kronesedler enn med 100 kroner sedler i køen, og billettøren hadde to 100 kronesedler i kassen forut for de m kundene?

Løsning: Sannsynligheten for at det oppstår veksleproblem er $4/5$.

Notér for hver kunde $+1$ for 100 kroneseddel og -1 for 200 kroneseddel, og la s_k være delsum etter kunde nr. k . Veksleproblem oppstår straks s_k blir negativ. Situasjonen kan illustreres med en graf med retning fra venstre mot høyre, der en ved hver node går opp ($+1$) eller ned (-1), så lenge det går an. Nivået i vertikal retning svarer til delsummene underveis. Utfallsrommet er alle mulige stier fra startnoden (til venstre) til sluttnoden (til høyre). Billettøren får et veksleproblem straks stien bryter den horisontale midtlinjen fra oversiden til undersiden. I grafen er tegnet et utfall uten veksleproblem med grønn strek ($+1+1-1-1+1+1-1-1$), og et utfall med veksleproblem med heltrukken rød strek ($+1+1-1-1-1+1+1-1$). Den stiplede fortsettelsen kommer vi tilbake til.



Vi antar at alle mulige stier er like sannsynlige og bruker regelen om «gunstige på mulige». Med $m=2n$ personer i køen blir antall mulige stier $\binom{2n}{n}$, som for $n=4$ er lik 70. Flertallet av stier går innom undersiden, og vi teller derfor de som ikke gjør det, og finner 14. Sannsynligheten for at det ikke oppstår noe veksleproblem er derfor $14/70=1/5$. I jakten på en generell formel får man for $n=2$ sannsynligheten $1/3$, og for $n=3$ sannsynligheten $1/4$. Er det virkelig slik at i det generelle tilfellet $n=2m$ er sannsynligheten for intet veksleproblem lik $1/(n+1)$ og dermed veksleproblem $n/(n+1)$? Ja, slik er det, og da bør det vel finnes et enkelt resonnement?

Her følger tre ulike resonnementer, ikke helt enkle, men interessante fordi de gjør bruk av helt ulike analytiske verktøy. Først tar vi et rent kombinatorisk resonnement.

La generelt $S(a,b)$ være en sekvens av a $+1$ 'ere og b -1 'ere. Antall slike sekvenser er

$B(a, b) = \binom{a+b}{a} = \frac{(a+b)!}{a!b!}$. Med $m=2n$ kunder i køen, har vi en $S(n, n)$ sekvens. En slik kø som leder til veksleproblem kaller vi $S^*(n, n)$ -sekvens. Disse har minst en negativ delsum, og er innom undersiden i figuren. Det er klart at antall mulige stier er $B(n, n)$, og vi skal vise at antall stier som gir veksleproblemer blir $B(n-1, n+1)$. Dette gir

$$\text{gunstige/mulige} = B(n-1, n+1)/B(n, n) = n/(n+1).$$

Her er et enkelt resonnement for telleren, som omgår omfattende analytisk regning:

Betrakt en vilkårlig $S^*(n, n)$ -sekvens (problemkø) og anta at veksleproblemet oppsto først ved kunde nr. k (som må være et odde-tall). Blant de $k-1$ foran i køen må det ha vært like mange $+1$ som -1 . Med nr. k kom en ny -1 , og blant de etter nr. k i køen må det så være en $+1$ mer enn det er -1 . Hvis vi bytter fortegn på disse $2n-k$ kundene, sitter vi igjen med en entydig $S(n-1, n+1)$ -sekvens. Det er dette som illustreres med de stiplede fortsettelsen av de røde linjestykkene i figuren. Slike sekvenser har sluttnode på undersiden, svarende til sluttdelsum på -2 . På den annen side, dersom vi har en vilkårlig slik $S(n-1, n+1)$ -sekvens (for problemkø eller ikke), gis en entydig $S^*(n, n)$ -sekvens (problemkø). Det ser vi ved at en $S(n-1, n+1)$ -sekvens må ha delsum -1 før eller senere, første gang ved kunde nr. k . Det er da like mange $+1$ som -1 før kunde nr. k , og en -1 mer blant de etter nr. k i køen. Dersom vi bytter fortegn for disse, blir det like mange $+1$ som -1 i hele køen, som derfor har tilordnet en entydig bestemt $S^*(n, n)$ -sekvens. Antall $S^*(n, n)$ -sekvenser må derfor være lik antall $S(n-1, n+1)$ -sekvenser.

Merk at antallet problemstier og ikke-problemtier er henholdsvis

$$B_n = \frac{n}{n+1} \binom{2n}{n} \quad \text{og} \quad C_n = \frac{1}{n+1} \binom{2n}{n}$$

Her er $\{C_n, n = 0, 1, 2, 3, 4, \dots\}$ følgen av såkalte Catalanske tall, etter Eugène Catalan (1814-1894), som dukker opp i mange ulike problemstillinger, noen allerede på 1700-tallet (bl. a. hos Euler). Resonnementet ovenfor er blitt kalt André's refleksjonsprinsipp, etter Désiré André (1840-1918), som anvendte det i en annen sammenheng (Bertrand's ballot problem). Vi kjenner også lignende resonnementer i forbindelse med passeringssannsynligheter i Wiener-prosessen.

I situasjonen der billettøren har to 100 kronesedler i kassen istedenfor ingen, blir en generell beregning tilsynelatende svært komplisert, men i tilfellet der det er akkurat to flere kunder med 200 kroner enn 100 kroner, åpner det seg en enkel mulighet: Start med en $S(n+1, n+1)$ -kø uten 100 kronesedler i kassen. Skriv ned uttrykket for sannsynligheten for at dette er en problemkø ved å betinge med omsyn på de to første i køen. Det er da tre muligheter: $(-1, -1)$ ($-1, +1$) leder direkte problem, $(+1, -1)$ leder til $S(n, n)$ -kø uten 100 kronesedler i kassen, og $(+1, +1)$ leder til en $S(n-1, n+1)$ -kø med to 100 kronesedler i kassen. Dette gir en ligning, der alle sannsynlighetene for problemkø som inngår er kjent unntatt den siste, som nettopp er den det spørres etter. Innsetting og løsning gir sannsynligheten for problemkø lik $(n-1)/(n+2)$. For $n=4$ blir de to mulighetene like sannsynlige. Merk at de to ekstra 100-kronesedlene i kassen mer enn kompenserer for at det er to flere som trenger vekslings av sin 200 kroneseddel.

En annen tilnærming til problemet går ut på å etablere konvolusjonsformelen

$$C_{n+1} = \sum_{k=0}^n C_k \cdot C_{n-k} ; C_0 = 1$$

Denne gir en ligning for den tilhørende genererende funksjon. Rekkeutvikling av løsningen gir C_n . Det som trengs for å etablere ligningen er at enhver uproblematisk $S(n+1, n+1)$ -sekvens kan entydig «faktoriseres» på formen $(+1)S_1(k, k)(-1)S_2(n-k, n-k)$, for en eller annen $k=0, 1, \dots, n$, der $S_1(k, k)$ og

$S_2(n-k, n-k)$ begge er uproblematiske sekvenser, og der $S_1(0,0)$ og $S_2(0,0)$ betyr «tom sekvens». Eksempelvis faktorerises den grønne sti i figuren slik: $(+1+1-1-1+1+1-1-1) \rightarrow (+1)(+1-1)(-1)(+1+1-1-1)$.

En tredje tilnærming til problemet er å etablere en differensligning i to variable.

La $P(a, b)$ være sannsynligheten for at en kø med $a+b$ avvikles uten veksleproblem, der a kunder har 100 kroneseddel og b kunder bare har 200 kroneseddel. Vi er interessert i $P(n, n)$. Da har vi

$$P(a, b) = \frac{a}{a+b} \cdot P(a-1, b) + \frac{b}{a+b} \cdot P(a, b-1); a \geq b$$

med bibetingelsene $P(a, b) = 0$; $a < b$ og $P(a, 0) = 1$. Løsningene av denne ligningen er (sjekk ved innsetting)

$$P(a, b) = \frac{a-b+1}{a+1}; a \geq b.$$

Herav følger som ventet at $P(n, n) = \frac{1}{n+1}$. Merk at denne løsningen egentlig innebærer at vi, ut fra antakelsen om at alle rekkefølger i utgangspunktet er like sannsynlige, bruker at etter hver ekspedert kunde, så er sannsynligheten for neste kundes pengeseddel bestemt av andelen gjenværende kunder med den gitte seddeltype. Det er selvsagt ganske opplagt. Likevel, gjør man «nye» antakelser underveis, bør man sjekke og påpeke at disse er forenlige med de antakelser man eventuelt gjorde i utgangspunktet. Mange lærebøker slurver med dette i sine eksempler.

møter og konferanser

Oppstart NSF Stavanger

Stein Andreas Bethuelssen | Stavanger



Anastasia Ushakova holder sitt foredrag

Høsten 2019 ble den første lokalgruppen til NSF opprettet i Stavanger. I motsetning til lokalforeningene så er NSF Stavanger underlagt NSF og har ikke egne statutter eller årsmøter. Gruppens formål er i første rekke å støtte opp om NSF's arbeid for statistikere og statistikkinteresserte i Stavangerregionen, blant annet ved å arrangere møter rettet inn mot medlemmer av NSF og andre statistikkinteresserte.

En tidlig høstkveld, onsdag 27. november 2019, avholdt NSF Stavanger sitt første (og siste?) oppstartsmøte. Møte fant sted på lunsjrommet til instituttet for matematikk og fysikk ved UiS og var en stor suksess. Over 20 statistikkglade deltok på møtet, hvorav ca halvparten ansatt ved UiS.

Det faglige programmet bestod av fire foredrag av lokale statistikkentusiaster. Først ut var Tore Selland Kleppe (UiS) som fortalte om sine interesser innen Monte Carlo metoder. Deretter fortalte Anastasia Ushakova (SUS) om hennes mangfoldige bakgrunn blant annet fra tiden som PhD ved NTNU, og nå i senere tid som statistiker ved Universitetssykehuset i Stavanger. Neste ut var Kaouther Hadji (Accenture) som fortalte om modellering innen biologi og finans ved bruk av stokastisk

analyse og Monte Carlo metoder, basert på hennes arbeid fra tiden som PhD ved Paris 6 og Postdoktor ved universitetet i Lyon. Det faglige programmet ble avsluttet med et foredrag av Jørgen Bølstad (UiS) som fortalte om hans bruk av Monte Carlo metoder og bayesiansk statistikk innen sosialvitenskap.

Det faglige programmet ble supplert med enkel bevertning og jevnt over sosial hygge. Gjengangsmelodien var at dette er noe som absolutt bør gjentas og videreutvikles! Det tas sikte på å avholde lignende tilstelninger minst en gang hvert semester.

Medlemsmøte NSF Trondheim



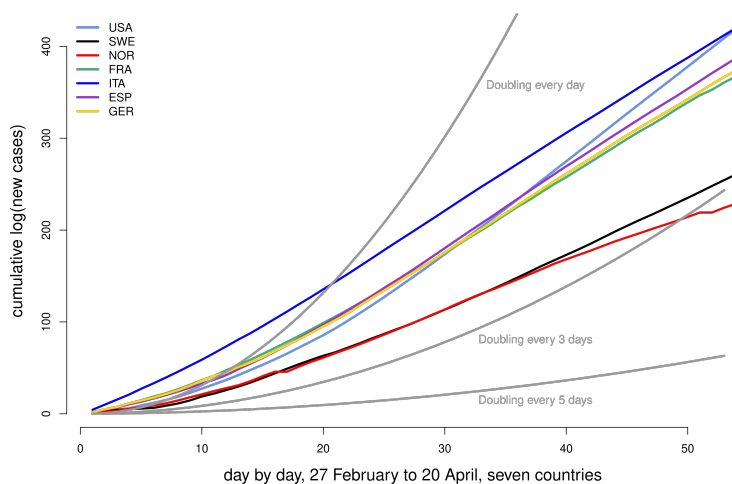
Den 9. mars hadde vi medlemsmøte i Norsk statistisk forening, avdeling Trondheim, i Lunsjrommet på Institutt for matematiske fag. Det deltok omtrent 30 medlemmer fra hele Trondheim. Se https://wiki.math.ntnu.no/tsf/mote_9_mars_20 for mer.

Monitoring the Level and the Slope of the Corona

Nils Lid Hjort og Emil Aas Stoltenberg | Oslo

Merk; denne artikkelen ble først publisert som en blog-post på [The FocuStat Blog](#).

The corona is mightily upon us, with different governmental decisions having drastic consequences, but in uncharted territory. Which country is wiser, Norway or Sweden? Have Italy and Spain passed their peaks? Here we study the two crucial data series of (1) new registered covid-19 cases and (2) new covid-19 related deaths, day by day, country by country. Via log-linear modelling we build practical statistical monitoring tools for assessing both the level and the slope of these two corona processes.



How life looks like, in 2020.

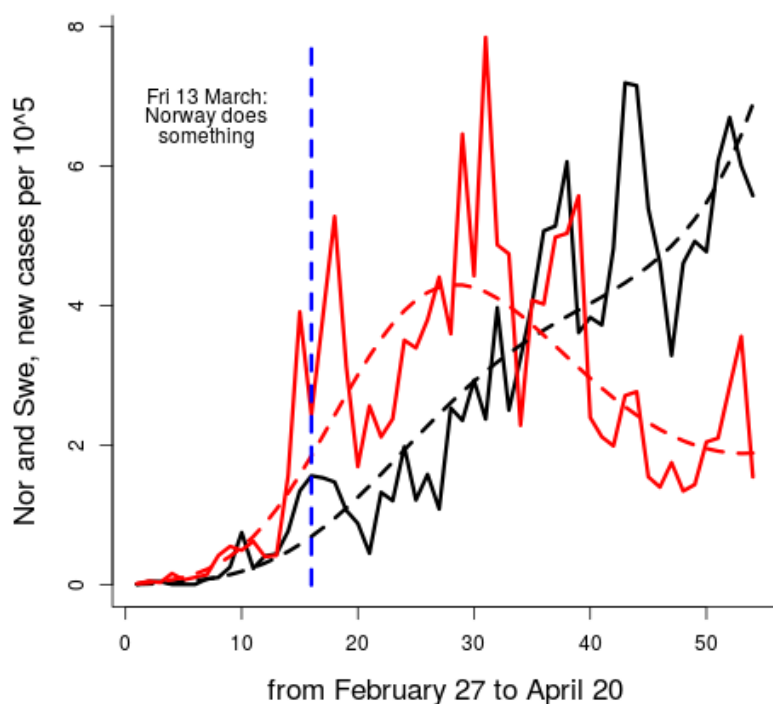
First-derivatives don't often make daily headlines, but these are strange times, and **flattening of curves** is on everybody's mind. In this blog post we build practical statistical monitoring tools for assessing the level and the slope of two corona-virus processes: the daily number of **new registered cases** of covid-19, and the daily number of **covid-19 related deaths**. Estimating first-derivatives comes, as with all statistical estimates, with uncertainty attached, and this blog post introduces nice uncertainty summaries regarding these how-steep-how-flat estimates.

A note of caution: We analyse the raw data as published by the **European Centre for Disease Prevention and Control**, and to what extent the two time series we use in this post are reflective

of the unknown registered-plus-unregistered-covid-19-cases times series, we do not attempt to answer in this blog post.

We're going to build statistical models for these series of N_t , number of new cases on day $t = 0, 1, 2, \dots$, clearly with one model and assessment analysis for the first data series (new registered cases) and another for the second series (new deaths). For simplicity we refer to these as [Type One](#) and [Type Two](#) data, respectively.

The full dataset we've used for this FocuStat Blog Post story, for the seven countries Sweden, Norway, France, Italy, Germany, Spain, the USA, which we've put together via the website pointed to, can be found [here](#), as a 54×14 data matrix. It has both [Type One](#) and [Type Two](#) data, with the clock set to day zero at February 27, $t = 0$, running onwards to April 20, $t = 53$, which happens to be today, the date of us writing up this report.



[Figure A](#): daily new registered cases, per 100,000, from February 27 to April 20, with Norway in red and Sweden in black. The dashed smooth curves are estimated trends from the fitted third-order log-linear model, and the vertical dashed blue line marks Friday March 13, where Norway decided to close down schools and universities and with yet further restrictions on travel and meetings.

This makes it possible for us to present analyses, assessments, and summary graphs of type [Figure A](#), with more to come below. This figure presents first of all the raw data, for the case of [Type One](#), for Norway and Sweden, over the time-window in question. It also shows estimated trend functions, which we come back to below. In addition

to yielding well-fitted curves from well-working models, the modelling apparatus allows us to reach inferential statements, clear comparisons between different countries (the validity of which, of course, depends on the raw data being comparable between countries), and to produce at least valid shorter-term predictions, say for the coming two weeks.

Since we're interested in comparisons with several different countries we're also normalising the count data, taking population sizes into account; specifically, our count series, for both [Type One](#) and [Type Two](#), are taken as [per 100,000 people](#). Whether this is an adequate measure for comparing countries is not clear, and it can be argued that one ought to consider smaller geographical entities than countries, for example cities (Oslo and Finnmark might exhibit rather different corona behaviour). Our models and methods would work well also for such smaller entities, with modest modifications. At any rate, comparison on the per 100,000 scale for Sweden and Norway appears reasonable.

At the time of writing this, 20 April, the corona processes have been taking their dramatic tours and tolls, in country after country, for some two months. In several countries the curves have managed to flatten out and even to display mild decrease, thanks to various drastic governmental interventions. So it might be a good time to assess both [levels](#) and [slopes](#), for evaluation, the next level of planning, and nowcasting.

The data and a class of statistical models

We should state a [couple of caveats](#) up front – first, we're taking the [Type One](#) data on board as they are, in spite of certain issues. These registration data partly and unavoidably reflect the testing regimes, and there are certain error rates at work. The extent to which the registration data reflect the registered + unregistered data is unknown, and, to complicate matters, probably differs from country to country.

The [Type Two](#) data are more drastically hard-core (a death is a death), but even with these there are debates regarding precisely what deaths are corona-decided and which might not be; see Kristiansen's [editorial in forskning.no](#) (April 19).

At any rate, even though some of these issues might be worked with further, using more advanced statistical modelling on top of the directly observed data series, the first statistical ambition is to understand and analyse the data we have, which in this blog post translates to estimating and assessing levels and slopes of the [Type One](#) and [Type Two](#) data. Good models and methods for these tasks also open the door to nowcasting (short-term predictions) and also for comparisons between countries and their strategies.

It turns out that essentially similar models and methods work for both [Type One](#) and [Type Two](#) data, though obviously

with different sets of parameters, estimates, prognoses, etc. But in a statistical methodology sense, work we carry out for modelling, handling, interpreting data of one type of $N_1, N_2, N_3, \dots$ can be transferred and applied to the other type. This is not to be taken for granted in advance, but is among our findings. Also, data on new deaths this week (Type Two) relate by definition to data on registered cases a few weeks back (Type One).

Though various extensions, modifications, and sophistications are possible here, both regarding the models and how to analyse them, our basic class of models takes the log-counts as a polynomial trend over time plus additive noise. The third-degree version of this has

$$Y_t = \log N_t = \beta_0 + \beta_1 x_t + \beta_2 x_t^2 + \beta_3 x_t^3 + \varepsilon_t,$$

for time points $x_t = t = 0, 1, 2, \dots$, with the ε_t being zero-mean error terms. We think about

$$m(t) = \beta_0 + \beta_1 x_t + \beta_2 x_t^2 + \beta_3 x_t^3$$

as the trend function, unfolding over time, typically with a drastic escalation from mid March before levelling off at mid April or some weeks later, before they are expected to point down again. We shall take special interest in [the present level](#) $m(t)$ and [the present slope](#) or derivative $m'(t)$, the latter a bit harder to assess well than the first.

The model thus described turns out to be a quite good one for the [Type One](#) data, even with the further statistical modelling simplification to take the error terms as independent with the same standard deviation, say σ . The degree 3 is arrived at via some model checking and model selection work, involving [the AIC and the FIC](#).

Of course [Type One](#) and [Type Two](#) data are different, in many regards. [The first](#) prominent feature is that the death numbers, of course, are a significant factor smaller in size; as of April 20, there are some 2,333,000 diagnosed corona cases world-wide and some 161,000 corona-defined deaths. This amounts to corona deaths being about 6.9 percent of corona diagnosed cases.

This is luckily [not](#) to be interpreted as saying that 6.9 percent of corona cases lead to death; testing volume is increasing, also in the presumed healthy population. At the same time, medical treatment gets better, in many countries. A [study in the Lancet](#) estimates the mortality rate of the coronavirus to about 3 percent, and a similar number was reported by the WHO in early March. The seasonal flu, by comparison, has a death rate of about 1 percent. For how the Covid-19 virus compares to other causes of death, see [this summary](#) in the New Atlantis, April 13.

[Secondly](#), the death count numbers display a bit more statistical irregularity than the numbers diagnosed. In statistical

modelling terms we're sometimes led to 4th order trend models, not 3rd; the residual variance is less stable; and for some countries there's further autocorrelation present (these are statements concerning the data series as per April 20; matters may stabilise over the coming few weeks).

Part of the point in our brief story is that modelling, assessments, and the building of good monitoring tools is easier and clearer on the logarithmic scale than on the original count-scale of things. Monitoring which deviations from the trend are inside normal or not becomes a clearer statistical job, as does the estimation of both level and slope, as we see below.

To showcase our models and assessment methods we choose as primary illustration the comparison between the Kingdoms of Norway and Sweden, regarding both [Type One](#) and [Type Two](#) data. Our models are in the category of «not too complicated but quite effective» for their intended purposes, which is to monitor the level and the slope. As mentioned above we're also analysing similar data series for other countries, however, and comment briefly on this below.

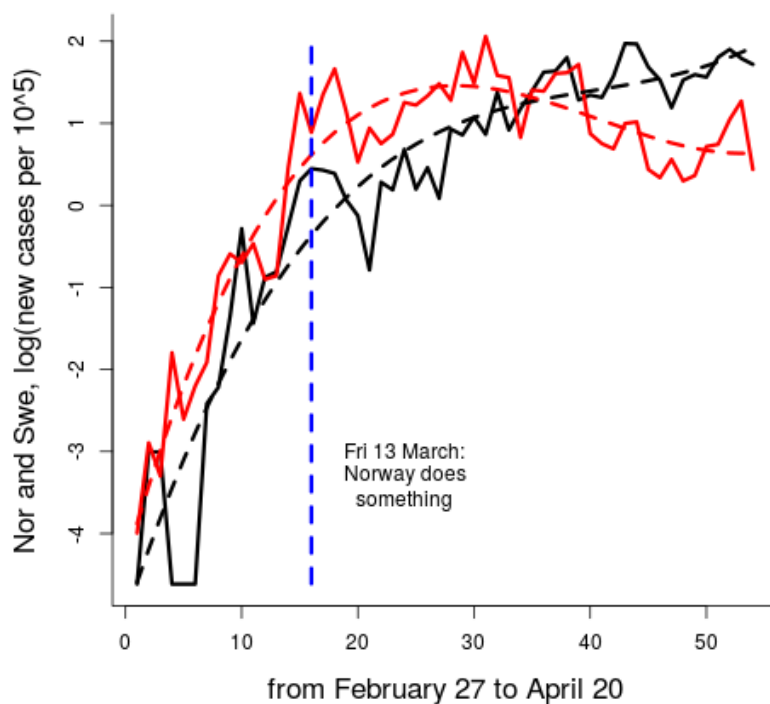


Figure B: the same statistical information as with [Figure A](#), but now presented on the log-scale, where it is easier to monitor and measure trends, deviations, and shifts. Norway's action day is again marked with the blue vertical line.

The new cases data series

Consider the corona-registration [Type One](#) data, the actual daily wiggly official up-and-down counts, for Norway (red) and Sweden (black), plotted in [Figure A](#) above. The vertical

dashed blue line marks the Day of Intervention for Norway, Friday March 13, where schools and universities were closed down, along with various other restrictions on travel, etc.

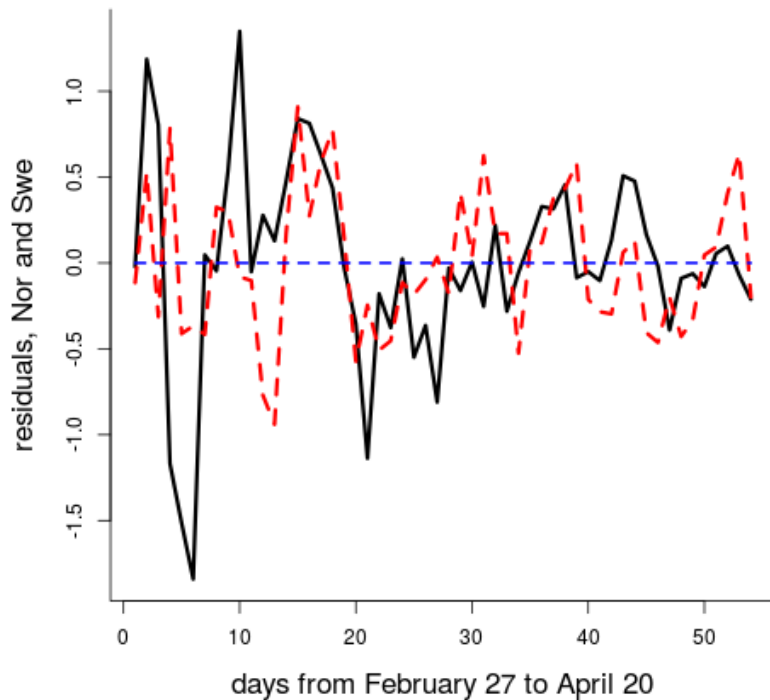
The dashed smoother lines, in both [Figures A and B](#), are trend estimates coming from our log-linear statistical models, briefly described above. Indeed, on the log-scale of [Figure B](#), the smooth trend functions is

$$\hat{m}(t) = \hat{\beta}_0 + \hat{\beta}_1 x_t + \hat{\beta}_2 x_t^2 + \hat{\beta}_3 x_t^3$$

which is brought back to original count-scale $\exp(\hat{m}(t))$ in [Figure A](#).

The variance level for the residuals has reached a reasonable equilibrium, see [Figure C](#), and there is even approximate normality, though this is not particularly important for our analyses.

There are various somewhat sophisticated techniques for estimating these coefficients, including weighted log-likelihood procedures with more weight given to the more recent observations, but even a straightforward least sum of squares method will work well here.



[Figure C](#): the residuals, from February 27 to April 20, for the Norwegian (red) and Swedish (black) [Type One](#) datasets, via the third-order log-linear model. The data series have reached a reasonable equilibrium, with the standard deviation level being stable. Also, the autocorrelation level is low.

The level and the acceleration

It's now time to [focus](#). Two parameters of primary interest are the level and the slope, and these quantities need precise enough definitions so that we can start working with them inside our log-linear model. There are indeed different nuances and variations here, but at least a starting point is to take

$$\phi_A = \frac{1}{k} \sum_{t \in I_k} (\beta_0 + \beta_1 x_t + \beta_2 x_t^2 + \beta_3 x_t^3)$$

for the level parameter, i.e. the model mean parameter averaged over the last k timepoints. Somewhat more involved, but similarly, we may take

$$\phi_B = \frac{1}{k} \sum_{t \in I_k} (\beta_1 + 2\beta_2 x_t + 3\beta_3 x_t^2)$$

as the operative slope parameter, the averaged derivative over the most recent timepoints. Having properly decided on precise focus parameters, we can go on to the technical business of estimating them, along with clear measures of uncertainty, so that relevant monitoring and testing can be carried out. We may check

- whether the slope is bigger for Sweden than for Norway (it is, as of today, but it might change in three weeks time);
- whether most of the European countries are roughly in the same boat (they're not);
- the degree to which Norway on May 15 is different from Norway on April 15, etc.

In particular, for each focus parameter, from ϕ_A and ϕ_B here to others, we can construct clear confidence curves – see [Cunen and Hjort's FocuStat Blog post](#) on that theme.

Without going into too much of the technical details here, we note that both ϕ_A and ϕ_B may be expressed as $\phi = w^t \beta$, a linear combination of the four β_j coefficients. This leads to estimator $\hat{\phi} = w^t \hat{\beta}$, which is approximately normal (well, depending on our precise modelling assumptions) with a formula for its standard deviation τ and even for its full distribution. For several of these focus parameters all of this leads to confidence curves of the type

$$cc(\phi) = |1 - 2G_{df}((\phi - \hat{\phi})/\hat{\tau})|,$$

with G_{df} the cumulative distribution function for a t -distribution with $df = n - 4$ degrees of freedom, n being the length of the time series under investigation. The ϕ_A parameter is the current expected level of the $Y_t = \log N_t$ process, which also translates to current median level $\exp(\phi_A)$ on the cases-counting scale. The ϕ_B is the current slope, on the log-level, which means that

$$\phi_B^* = \exp(\phi_B)$$

is the slope parameter on the counting scale. Confidence curves shown in [Figures D and F](#) are shown for this natural-scale acceleration parameter $\exp(\phi_B)$ rather than for ϕ_B on the log-scale. Countries where $\exp(\phi_B)$ is a bit above 1.00 must hope that it soon enough is pushed down below 1.00.

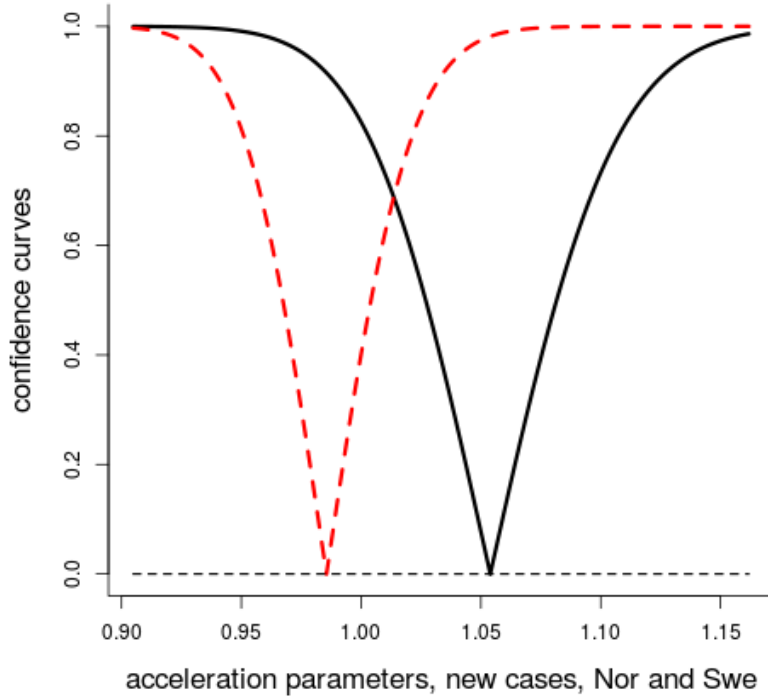


Figure D: confidence curves for the acceleration or slope parameters $\exp(\phi_B)$, for Norway (red) and Sweden (black), for the [Type One](#) new corona diagnosed cases data. For tomorrow, Sweden can expect 1.054 times today's number (5.4 percent more), whereas Norway can reckon with 0.984 times today's number (1.6 percent decrease). The difference between above 1.00 Sweden and below 1.00 Norway was significant a week ago, but has evened out, and is not significant as of April 20.

Norway and Sweden, the new cases data

In [Figure D](#) we show confidence curves for this acceleration like parameter $\phi_B^* = \exp(\phi_B)$, working on the original counting scale, for the [Type One](#) data, for Norway and Sweden, as of April 20. The acceleration estimates, on this scale, are 0.984 for Norway, indicating a downward trend, and 1.054 for Sweden, signalling a slight but not dramatic upper trend. Translated to estimates for tomorrow's numbers, this indicates that

$$N_{t+1} \doteq 0.984 N_t \text{ for Norway, } N_{t+1} \doteq 1.054 N_t \text{ for Sweden.}$$

If the situation stays stable, for say a week, this translates further to estimated factors $1.054^7 = 1.445$ times today's level for Sweden and $0.984^7 = 0.893$ times today's level for Norway; these are exponential mechanisms at work.

Surely also worth pointing to is that the two countries are at roughly the same level, after all (when numbers are normalised according to population size, as we've done in [Figures A and B](#)). Confidence curves for the level ϕ_A (not displayed here) show that the current Swedish level is about $\exp(1.72) = 5.81$, and the Norwegian considerably lower $\exp(0.66) = 1.94$, on the scale of new daily corona cases per 100,000; see [Figures A and B](#). The difference is significant, but might even out over the coming weeks and months.

It is relevant here to point out that the difference between acceleration numbers 0.984 and 1.054 for Norway and Sweden is not significant, as of April 20, see [Figure D](#); using data from one week ago, however, the difference was clear and significant. So the corona processes are dynamic, and initial differences might be evening out in the longer marathon run of things.

So there are certain differences between Norway and Sweden, as portrayed in [Figures A, B, D](#), but based on the raw data and our log-linear models they do not appear to be that big. Also, just one week of incoming new data, from April 13 to April 20, have made differences smaller (we've checked, using our models and methods). Whether the differences present can be attributed to the stricter measures imposed in Norway compared to in Sweden is a complicated question that we cannot answer with these data at the present time.

Also, the much harder question of the long term economic consequences of the differing strategies must be left for the thousands of Masters and PhD theses of statistics, economics, political science, epidemiology, etc. that surely will be produced in the coming years.

Death count statistics, for seven countries

For simplicity of presentation we limited attention above to the [Type One](#) data of counting new corona cases, comparing Norway and Sweden. As mentioned in the introduction, the same type of models and hence methods essentially work also for the hard-core [Type Two](#) data, counting corona-inflicted deaths day by day, country by country.

These data exhibit somewhat more irregular behaviour, however, calling in cases for a 4th order rather than a 3rd order polynomial trend function, and also for more carefully handling and modelling of standard deviation level, say σ_t at time t , and with autocorrelation.

We do not have room here for showing and discussing full analyses from our models and methods, but we have indeed properly analysed [Type One](#) and [Type Two](#) data for the seven countries Sweden, Norway, France, Italy, Germany, Spain, the USA, from February 27 to April 20. [Figure E](#) shows one type of output from such modelling work, with actual log-counts per 100,000 along with fitted trend functions, here for Norway,

Sweden, France, Italy. The fitted trends are reasonable, though with different levels of accuracy for different countries. Norway clearly has a lower death rate per population size. Intriguingly and positively, these four nations appear to have levelled out, regarding the death counts, and France and Italy in particular are on their ways downwards. Also Sweden is close to its envisaged peak, the flattened curve. More accurate information can be gleaned from our slope or acceleration analysis, below.

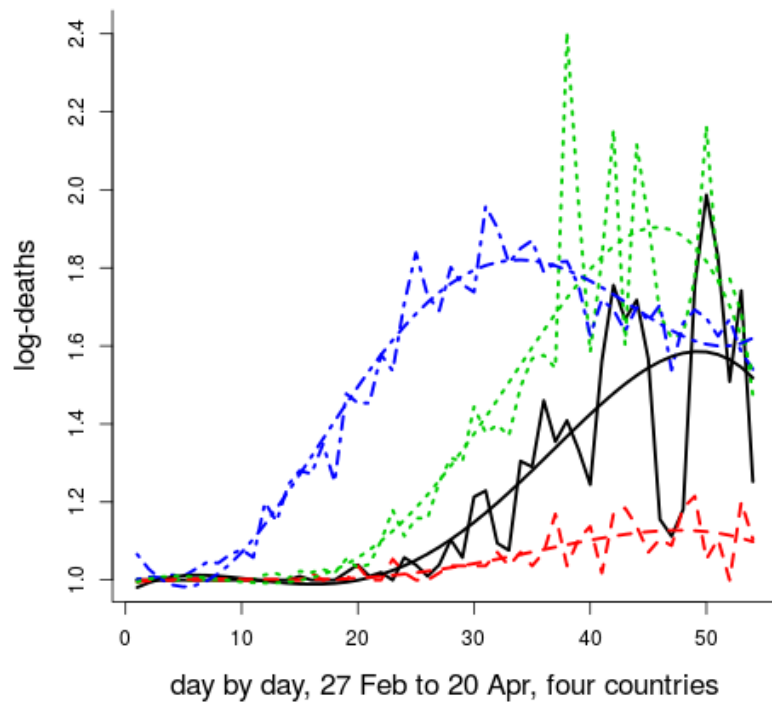


Figure E: log-count [Type Two](#) data, i.e. number of corona deaths, per 100,000, on a log-scale, for Norway (red), Sweden (black), France (green), Italy (blue). The irregular curves are the actual log-counts, with the smoother curves coming from our 4th order log-polynomial trend models.

For the more regular corona cases counting [Type One](#), analysis similar to that displayed in [Figure D](#), for the other five nations, show that the $\exp(\phi_B)$ parameters are not far from the stability value 1.00, with some nations just above 1.00 (like Sweden, an upward trend, still) and others just below (like Norway, the downward trend has started).

Also for the death counts [Type Two](#) data, these acceleration parameters are quite close to 1.00, as seen in [Figure F](#). France and Italy are the current best in class (again, as of April 20), with clearly shaped descents (acceleration less than 1.00), whereas the other five countries, including Norway and Sweden, are close to stability. We see, however, that the Swedish acceleration parameter is far less certain than the Norwegian one, reflecting more variable up-and-down deaths data on the Swedish side.

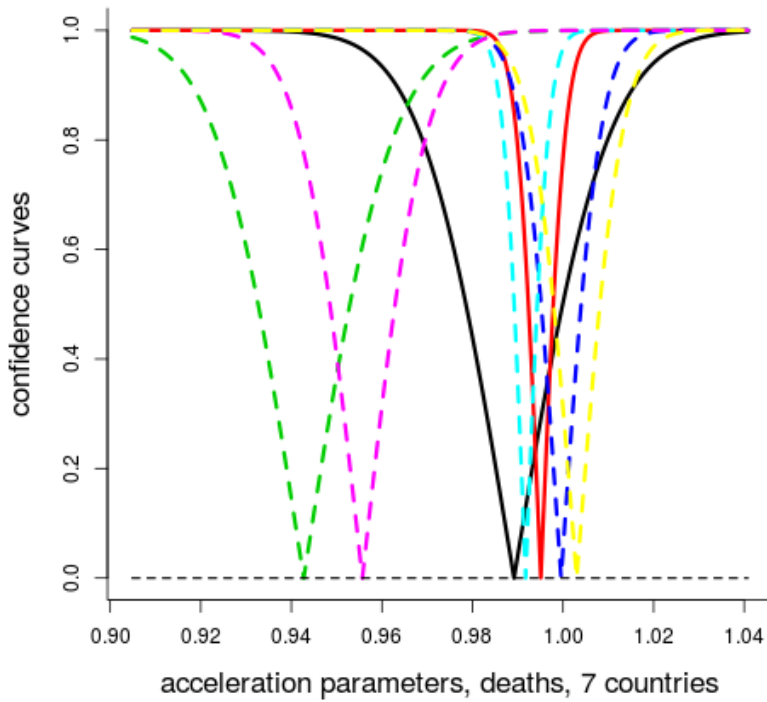


Figure F: confidence curves for the acceleration parameters $\phi_B^* = \exp(\phi_B)$, for the corona death **Type Two** data, for seven countries, as of April 20. The two to the left are France and Spain, with clear downward slopes, i.e. accelerations significantly lower than 1.00. The Norwegian and Swedish parameters are similar in size, actually close to the stability level 1.00, but the Swedish is currently more uncertain.

A few notes

A: Of course there is a big literature on time series which can be brought to use here. Most of the models and methods in that big field are however geared up for mostly stationary phenomena (if not stationary in the mean level, then perhaps stationary in the slope), whereas the dramatic corona data series currently unfolding in front of us display dynamic and nonstationary characteristics. It is in this spirit we've modelled these data series more or less from scratch. We trust the analysis tools we've constructed will be useful for the broader business of monitoring changes, estimating the near future, perhaps country by country. Machinery from Cunen, Hjort, Nygård (2020) can be used for addressing the change points, or regime shifts, and measuring their impacts.

B: Our models and methods may also be used for so-called nowcasting, or shorter-term forecasting. We may e.g. carry out simulations for the number of new corona deaths over the coming three week, country by country, where the statistical machinery involves both the fitted models and the precision of the implied parameters. We do not know precisely how other

statisticians in the corona field do this, like in the Jewell and Jewell articles from earlier in April, but we could attempt to match those in our later work.

C: It is difficult to predict, particularly about the future. Another difficult task, which is nevertheless tempting to work on, is the counterfactual one, the «what would have happened if» questions. How many US lives could have been saved, if the social distancing rules had been set in motion e.g. two weeks earlier? We refrain from analyses of this type at the moment (April 20), but the questions are valid, and our type of models can be used as playing and analysis ground.

D: Our statistical models and methods work well, as seen from various model diagnostics and model selection schemes. They could nevertheless be improved upon, with a more careful modelling of the typical shapes of the trend functions. Presumably theoretical or empirical studies would lead to parametrised versions $m(t, \theta)$. Such $m(t, \theta)$ could for example be informed by, perhaps even derived from, an underlying SIR-model or one of its more complex cousins. So in a sense we've just let the data speak for themselves, without reading theoretical papers about how the trend functions should look like under different sets of circumstances. With more efforts here we would also be able to sharpen tools of meta-analysis, how to analyse say a dozen countries jointly, as opposed to analysing each separately. The II-CC-FF setup for combining information across diverse sources should be useful here, see Cunen and Hjort (2020b).

E: From the tradition and viewpoints of classically oriented university-type statisticians it might be said, as a little meta-comment, that it is refreshingly unusual to work with contemporary data, even in the operative sense of finishing a report where the latest world-data to be put into our formulae and algorithms are from today. Hjort has worked with *medieval literature from the 1460ies, war and conflict data from 1824 to 2003, fisheries time series from 1859 to 2014, and Olympic history from 1924 to 2018*, so doing statistics on 20 April 2020 and writing up a full story on 20 April 2020 using data up to 20 April 2020 is a change of scale.

Thanks

We appreciate comments from and ongoing discussions with Ingrid Glad, Baard Meidell Johannesen, and Sidsel Kreyberg.

A few references

Claeskens, G., Hjort, N.L. (2008). Model Selection and Model Averaging. Cambridge University Press, Cambridge.

Cunen, C., Hermansen, G.H., Hjort, N.L. (2018). Confidence distributions for change-points and regime shifts. *Journal of Statistical Planning and Inference* vol. 195, 14-34.

Cunen, C., Hjort, N.L. (2020a). Confidence Curves for Dummies. *FocuStat Blog Post*, April 2020.

Cunen, C., Hjort, N.L. (2020b). Combining information across diverse sources: the II-CC-FF paradigm. *Scandinavian Journal of Statistics* [to appear].

Cunen, C., Hjort, N.L., Nygård, H.M. (2020). Statistical Sightings of Better Angels: analysing the distribution of battle-deaths over time. *Journal of Peace Research*, vol. 51, 1-16.

Hjort, N.L., Schweder, T. (2018). Confidence distributions and related themes. General introduction article to a Special Issue of the *Journal of Statistical Planning and Inference* dedicated to this topic, with eleven articles, and with Hjort and Schweder as guest editors; vol. 195, 1-13.

Hjort, N.L. (2020). Koronakrisen: plutselig ble statistikk allemannseie. Interview in *Titan*, University of Oslo.

Jewell, B.L., Jewell, N.P. (2020). The huge cost of waiting to contain the pandemic. Letter in *New York Times*, April 14.

Jewell, B.L., Lewnard, J.A., Jewell, B.L. (2020). Predictive mathematical models of the Covid-19 pandemic: underlying principles and value of projections. *JAMA Network*, April 16.

Kristiansen, N. (2020). Derfor kan vi ikke stole på tallene over døde og smittede av koronaviruset. Editorial, *forskning.no*, April 19.

Schweder, T., Hjort, N.L. (2016). *Confidence, Likelihood, Probability: Statistical Inference with Confidence Distributions*. Cambridge University Press, Cambridge.

Stoltenberg, E. Aa. (2020). How dangerous is the corona virus? University of Oslo Open Day talk, March 2020.

På leting etter sannheten bak kvantemekanikken

Inge Helland | Oslo

WARNING: Richard P. Feynmann said that attempting to understand quantum mechanics causes you to fall into a black hole, never to be heard from again.

Denne advarselen er hentet fra Richard Gill's hjemmeside. Richard lener sin egen tolkning av kvantemekanikken på arbeider av en matematisk fysiker som heter Slava Belavkin, men han har også flere referanser til andre fysikere.

Jeg har også lest mange arbeider av fysikere og registrert at de er fundamentalt uenige, så i utgangspunktet har jeg valgt ikke å lene meg til noen av dem. Det er en meget stor litteratur om grunnlaget for kvantemekanikken. Jeg har samlet i denne litteraturen, og prøvd å finne min egen tolkning ut fra et slikt nærmest tilfeldig, men likevel subjektivt utvalg. Resultatet ble en bok på Springer [1], og en rekke artikler, delvis med ufullstendige konklusjoner.

Er da boken fullstendig? Nei. Kapittel 4, der jeg prøver å finne et eget grunnlag, inneholder en feil, og er også nokså uferdig. I løpet av det siste året har jeg forsøkt å rette opp dette. En artikkel er publisert i Journal of Mathematical Physics, og har fått sitt erratum. Status nå er at jeg har sendt en mer omfattende artikkel til Foundations of Physics; selvfølgelig kan denne også vise seg å være ufullstendig.

Hva kan man konkludere ut fra alt dette? Vel, igjen må jeg gjøre et valg. Blant den store fysiker-litteraturen har jeg nå gjort et nytt lite, subjektivt utvalg, og skal prøve å forklare noen av grunnlagsproblemene ut fra disse artiklene.

Den mest uforståelige grunnlagssetningen er allerede denne: En tilstand av et fysisk system er en abstrakt vektor i et komplekst vektorrom (mer presist et Hilbertrom). Dette er utgangspunktet for dusinvis av lærebøker; en forholdsvis lesbar populær innføring er [2]. Jeg har prøvd å påpeke at disse vektorene under visse omstendigheter er i enentydig korrespondanse med 1) et fokusert spørsmål til naturen sammen med 2) et skarpt svar på dette spørsmålet [3], men har hittil fått forholdsvis liten respons fra fysikerne på dette utspillet. En yngre fysiker som heter Philip Höhn sendte meg imidlertid nylig en interessant melding om at han hadde bevist det samme i et viktig spesialtilfelle. Og nettopp sendte fysikeren Peter Morgan meg en artikkel om emnet som jeg skal studere nærmere.

La oss akseptere fysikernes utgangspunkt: En tilstand til et system er gitt ved en søylevektor ψ . La $\psi^\dagger$ være den komplekse konjugerte rekkevektoren. Sannsynlighetene i kvantemekanikken er da gitt ved Born's formel: Anta at systemet er preparert i tilstanden ψ_a . Sannsynligheten for at en ideell måling på systemet gir tilstanden ψ_b er $|\psi_b^\dagger \psi_a|^2$.

Denne formelen gjelder i utgangspunktet for systemer som har et endelig antall tilstander d , som enten kan beskrives ved en av de ortogonale vektorene $\psi_{a1}, \dots, \psi_{ad}$, eller ved en av de ortogonale vektorene $\psi_{b1}, \dots, \psi_{bd}$. Med en ideell måling mener jeg at vi ser bort fra usikkerhet i måleapparatet. En slik usikkerhet må etter min mening modeleres ved en statistisk modell.

Et enkelt eksempel er et elektron-spinn. En spinn-vektor har den spesielle egenskapen at komponenten av den i en vilkårlig retning a bare kan ha en av to verdier: Tilstanden ψ_{a-} gir verdien -1 , mens (den ortogonale) tilstanden ψ_{a+} gir verdien $+1$. Hilbert rommet er altså to-dimensjonalt. Det kan vises at tilstandsvektoren er entydig bestemt av valget av a , altså av spørsmålet: Hva er komponenten i retning a ? sammen med svaret -1 eller $+1$. En tilstand holder seg konstant inntil det tidspunktet hvor det eventuelt gjøres en måling.

Hvordan skal vi tolke en kvantetilstand? Noen mener den beskriver virkeligheten slik den er (ontologisk tolkning), andre mener at den beskriver vår (en observatør's) kunnskap om virkeligheten (epistemisk tolkning). Av flere grunner heller jeg mot det siste.

Flere diskusjoner dreier seg om to observatører, kalt Alice og Bob, som befinner seg så langt fra hverandre at de ikke kan kommunisere. Midt mellom disse er det en kilde til det som kalles to sammenfiltrede partikler; é partikkel sendes mot Alice og én mot Bob. Disse partiklene har ifølge kvantemekanikken en merkelig egenskap: Hvis Alice måler spinnkomponenten i en retning a og får verdien $+1$ (-1), vil verdien av spinnkomponenten i retning a for Bob's partikkel automatisk bli -1 ($+1$). Det spesielle er at dette gjelder uansett valg av a .

Dette er utgangspunktet for (Bohm's versjon) av den berømte debatten mellom Albert Einstein [4] og Niels Bohr [5], en debatt som fremdeles er like het [6]. Einstein kalte fenomenet 'a spooky action at a distance', og tok det som tegn på at kvantemekanikken er ukomplett. Han brukte resten av sitt liv på å lete etter en mer komplett teori. Bohr var mer pragmatisk. De fleste fysikere er nok enige med ham i dag, men det er også en almen følelse av at Bohr var noe uklar på hva han egentlig mente.

En ny vri på denne debatten kom ved John Bell [7, 8]. Bell så på en mer generell situasjon: Alice har valget mellom å måle spinn komponenten i en av to retninger a og b , mens Bob har valget mellom to retninger c og d . De teoretiske responsene kalles henholdsvis $\theta^a, \theta^b, \eta^c$ og η^d , som altså alle tar verdiene ± 1 . Av den siste grunnen har vi nødvendigvis at ett av tallene $\theta^a + \theta^b$ og $\theta^b - \theta^a$ har tallverdi 2, mens det andre er 0, og det følger at:

$$(\theta^a + \theta^b)\eta^c + (\theta^b - \theta^a)\eta^d = \pm 2,$$

og derfor spesielt

$$\theta^a \eta^c + \theta^b \eta^c + \theta^b \eta^d - \theta^a \eta^d \leq 2. \quad (1)$$

Anta nå at vi gjør en rekke målinger, der enten a eller b hver gang velges på en eller annen måte av Alice, mens c eller d på en eller annen måte velges av Bob. Vi får da estimerer $\hat{\theta}^a, \hat{\theta}^b, \hat{\eta}^c$ og $\hat{\eta}^d$, som vi igjen kan anta alle tar verdiene ± 1 . Gjør nå en avgjørende antakelse: Se på analogen til (1), og anta at produktene her har mening som tilfeldige variable, slik at vi kan ta forventning av hvert ledd i summen. I den fysiske litteraturen kalles dette en antakelse om lokal realisme, og det gir formelt ulikheten:

$$E(\hat{\theta}^a \hat{\eta}^c) + E(\hat{\theta}^b \hat{\eta}^c) + E(\hat{\theta}^b \hat{\eta}^d) - E(\hat{\theta}^a \hat{\eta}^d) \leq 2. \quad (2)$$

Dette er en av Bell's ulikheter (CHSH ulikheten), og den har spilt en stor rolle i flere ti-år med debatt. Poenget er først at man kan velge a, b, c og d slik at ulikheten ikke holder sett fra et kvantemekanisk perspektiv. Betyr dette at det er noe galt med kvantemekanikken, eller er det resonnementet over som ikke holder?

Det neste spørsmålet er følgende: Eksperimentet over kan gjøres helt konkret, enten med spinn av partikler, eller - som er helt analogt - polarisasjon i forskjellige retninger av fotoner. Hvis man da velger settinger a, b, c og d der kvantemekanikken viser brudd på Bell's ulikhet, vil eksperimenter vise brudd? De første eksperimentene av denne typen ble utført av Aspect og medarbeidere [9]. Det kom snart en rekke innvendinger og krav som slike eksperimenter måtte tilfredsstillte for at de skulle være overbevisende. De endelige eksperimentene som tilfredsstilte alle krav ble først utført i 2015 [10,11]; konklusjonen var meget entydig og skapte stor oppmerksomhet blant fysikerne: *Det finnes eksperimentelle situasjoner der Bell's ulikheter klart er brutt*. Dette er samtidig et argument for at prediksjoner fra kvantemekanikken ser ut til å holde i praksis.

Det er altså noe som ikke stemmer ved overgangen fra analogen av (1) til (2). Fra et statistisk synspunkt, hva er galt? Etter min mening kan det være litt belysende med følgende resonnement: Vi har en serie med data om retninger og responser fra begge observatørene. Disse kan modelleres ved en statistisk modell. Retningene valgt av Alice er samlet i en stokastisk variabel Z . Denne er ansillær, og fra betingingsprinsippet i statistikk skal vi betinge den statistiske analysen på Z . Tilsvarende skal den statistiske analysen av Bob's data betinges med hensyn på hans Z . Derfor må all forventning være knyttet til en observatør, enten Alice eller Bob.

Anta at vi konsentrerer oss om Alice, og at hun i en bestemt sett av kjøring velger retning a og får en verdi for $\hat{\theta}^a$. Ut fra Born's formel kan hun da finne tilnærmete sannsynlighetsfordelinger for $\hat{\eta}^c$ og $\hat{\eta}^d$, og dermed beregne

første og siste ledd i (2). Men hun har ingen mulighet til å beregne det andre og tredje leddet ut fra disse kjøringene. Merk at vi hele tiden betinger med hensyn til valg av retning. Tilsvarende blir det to ukjente ledd om hun velger retning b , og tilsvarende har Bob ingen mulighet til å beregne venstre side av (2).

Dette er situasjonen om vi ser på forventninger rett etter at målingene er gjort. Hva om Alice og Bob møtes senere og utveksler data? Det var først og fremst denne situasjonen Bell tenkte på med sine ulikheter, og det er denne situasjonen som svarer til eksperimentene som ble gjort i 2015. Diskusjonen om lokal realisme blant fysikerne de siste årene har skutt fart. Noen mener at naturen i en eller annen forstand må være ikke-lokal, men dette strider mot den spesielle relativitetsteorien. Flere og flere mener nå at en må gi slipp på antakelsen om streng realisme i naturen. Dette er vanskelig å fatte, og mange stritter fremdeles imot. Enkelte forskere har brukt innsatsen sin på å finne huller i Bell's argumentasjon. Richard Gill har brukt mye tid og energi helt til det siste på å argumentere imot slike forsøk. Bell hadde helt rett i sine resonnementer, mener han. Vi må spesielt gi opp forsøkene på å innføre streng realisme igjen. Carlo Rovelli, som jeg nevner nedenfor, har også argumenter mot en realistisk tolkning av kvantemekanikken.

En kan si at det enkle resonnementet mitt over med bruk av betingingsprinsippet bygger på en epistemisk tolkning både av kvantesannsynlighetene og sannsynligheter bak modeller for eksperimenter. Det er imidlertid mye som kan diskuteres her. Jeg sendte nylig et diskusjonsinnlegg til en nett-diskusjons-gruppe grunnlagt av Richard, med en argumentkjede som var beslektet med resonnementet over. Dette førte til en interessant debatt som jeg selv også har lært endel av.

Når det gjelder betingingsprinsippet hadde jeg i 1995 en artikkel i *American Statistician*, en artikkel som også inneholdt et 'moteksempel'. Dette førte til reaksjoner blant annet fra James Berger, Murray Aitkin og Dennis Lindley. Skulle jeg ha svart disse i dag, hadde jeg nok hatt litt mer respekt for autoritetene. Et tentativt forslag til en ny formulering av betingingsprinsippet finnes i [1]. Jeg skulle gjerne hatt reaksjoner på dette forslaget.

Etter min mening kan en fra et epistemisk synspunkt skissere forklaringer på en rekke såkalte paradokser i kvantemekanikken som Schrödinger's katt, Wigner's venn og tospalte-eksperimentet. Jeg har også en utledning av Born's formel ut fra visse forutsetninger i [1].

Når det gjelder kvantemekanikkens begrep om tilstander som komplekse vektorer, har jeg nå prøvd å nærme meg dette ved å bruke gruppeteori og grupperepresentasjonsteori. Dette er ikke-trivielt, men det er oppmuntrende at de samme matematiske teknikkene brukes i grunnlaget for kvantefeltteori

og fysikernes standardmodell [12].

En meget interessant observasjon er at kvante-modeller nå ikke bare brukes på mikroskopiske partikler, men også på psykologiske fenomener og i tilknytning til desisjonsteori [13].

Siden fysikerne er så uenige om grunnlaget for kvantemekanikken, skal en da stille seg skeptisk til hele den fysiske litteraturen? Nei, det mener jeg ikke. Etter mye leting synes jeg at Carlo Rovelli's 'Relational Quantum Mechanics' [14] er et godt utgangspunkt, også for en grundigere diskusjon [15] av Einstein-Podolsky-Rosen eksperimentet.

Det har skjedd mye siden Richard Feynman skrev sin advarsel. Blant annet har Stephen Hawking funnet ut [16] - ved å bruke kvanteteori - at informasjon kan unnslippe fra svarte hull.

Det er virkelig fint å ha en hobby nå i disse coronna-tidene som føles meningsfull. For meg er denne forskningen for lengst kommet så langt at det også har fått konsekvenser for mitt syn på endel områder innenfor filosofi (et notat om dette kan fås av meg) og innenfor religion [17]. Det siste notatet er lastet ned av mer enn 900 personer ifølge ResearchGate, så det har nok fylt et behov.

Men hele problemstillingen koker ned til følgende, slik jeg ser det: Hvilke autoriteter kan vi feste sin lit til? Kan vi stole på alle fysikernes resonnementer, spesielt de som ender med uenighet, eller skal vi prøve å være uavhengige og se det hele utenfra, for eksempel ved å insistere på å argumentere som statistikere eller matematikere? Jeg har valgt det siste, og vil derfor spesielt fortsette med å identifisere meg med kolleger fra min egen fag-disiplin. Og samtidig studere argumenter innenfor grunnlagsforskningen i statistikk. Men jeg kan ta feil.

Alle kommentarer til betraktningene over vil være velkomne.

Referanser

[1] Helland, I.S. (2018). *Epistemic Processes. A Basis for Statistics and Quantum Theory*. Springer. New York.

[2] Susskind, L. and Friedman, A. (2014). *Quantum Mechanics. The Theoretical Minimum*. Basic Books, New York.

[3] Helland, I.S. (2019). When is a set of questions to nature together with sharp answers to those questions in one-to-one correspondence with a set of quantum states? arXiv: 1909.08834 [quant-ph].

[4] Einstein, A. Podolsky, B. and Rosen, N. (1935). Can quantum-mechanical description of physical reality be considered complete? *Phys. Rev.* 47, 777.

[5] Bohr, N. (1935). Can quantum-mechanical description of physical reality be considered complete? *Phys. Rev.* 48, 696.

[6] Kupczynski, M. (2017). Can we close the Bohr-Einstein debate? *Phil. Trans. R. Soc. A Math. Eng. Sci.* 375 (2106): 20160392.

[7] Bell, J.S. (1964). On the Einstein-Podolsky-Rosen paradox.

Physics 1, 195.

[8] Bell, J.S. (2004). *Speakable and Unspeakable in Quantum Mechanics*. Cambridge UP, Cambridge.

[9] Aspect, A., Grangier, P. and Roger, G. (1982). Experimental test of Bells inequalities using time-varying analyzers. *Phys. Rev. Lett.* 49, 1804.

[10] Giustina, M. et al. (2015), Significant-loophole-free test of Bell's theorem with entangled photons. *Phys. Rev. Lett.* 115, 250401.

[11] Shalm, L.K. et al. (2015). Strong loophole-free test of local realism. *Phys. Rev. Lett.* 115, 250402.

[12] Robinson, M. (2011). *Symmetry and the Standard Model*. Springer, New York.

[13] Busemeyer, J.R. and Bruza, P.D. (2012). *Quantum Models of Cognition and Decision*. Cambridge UP, Cambridge.

[14] Stanford Encyclopedia of Philosophy (2019). Relational quantum mechanics. <https://plato.stanford.edu/entries/qm-relational/>

[15] Smerlak, M. and Rovelli, C. (2007). Relational EPR. *Found. Phys.* 37, 427.

[16] Hawking, S. (1993). *Black Holes and Baby Universes and Other Essays*. Bantam Dell Publishing Group, New York.

[17] Helland, I.S. (2017). The conception of God as seen from research on the foundation of quantum theory. *Dialogo* 4 (1), 259.

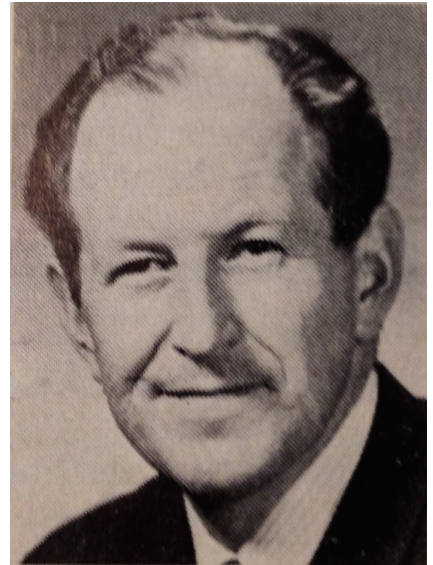
Tilbakeblikk: Den første nordiske konferanse i matematisk statistikk, Aarhus 1965

Jostein Lillestøl | Bergen

Initiativet til de nordiske konferansene i matematisk statistikk ble tatt av Arnljot Høyland og Ole Barndorff-Nielsen, som møttes under sine studier i California i 1962-64, Høyland ved Berkeley og Barndorff-Nielsen ved Stanford. Møteplassen var Berkeley-Stanford seminarene, der to sterke fagmiljøer i matematisk statistikk møttes regelmessig for utveksling av ideer. På den tiden var fagmiljøet i Norden lite, i Norge fantes matematisk-statistikere bare ved Universitetet i Oslo, der Høyland da var universitetslektor.

De to var enige om at det var behov for mer faglig kontakt over de nordiske landegrensene, og tanken om en nordisk konferanse lå nær. Behovet for møter der mange forskere la fram sin siste forskning i korte foredrag var allerede dekket gjennom de store internasjonale konferansene. Isteden valgte man et konsept med tre hovedforedrag, hver på tre timer, gitt av merittert forsker på det aktuelle fagfelt, som ga en oversikt over feltet, med utsyn til forskningsfronten. Etter foredragene var det avsatt rikelig tid til diskusjon, typisk to timer til hvert foredrag, med innleder som var godt forberedt. Støtte fra sentrale fagpersoner i de nordiske land ble sikret, i første rekke fra Erling Sverdrup i Norge, Anders Hald i Danmark, Herman Wold i Sverige og Gustav Elfving i Finland. Planleggingen skjedde utenom de statistiske foreninger i de respektive land. På denne tiden var den norske foreningen fremdeles en Oslo-klubb, med tyngdepunkt i Statistisk Sentralbyrå. Det statistiske miljøet ved Universitetet i Aarhus, med Barndorff-Nielsen i spissen, sa seg villig til å være vertskap for den første konferansen, som fant sted i juni 1965. De tre temaene på konferansen møtet var «Neyman Pearson teoriens stilling i dag» (Erling Sverdrup, Norge), «Bayesianske metoder» (Gustav Elfving, Finland) og «Item-analyse» (Georg Rasch, Danmark).

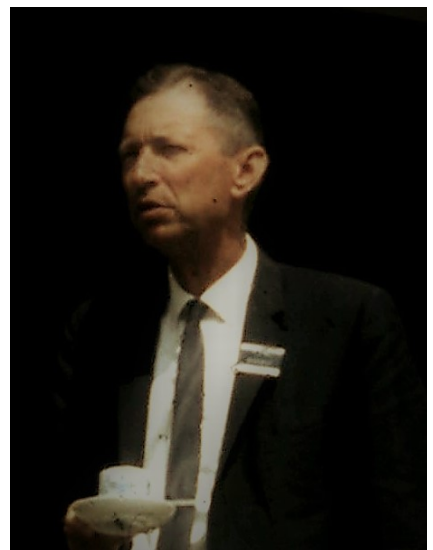
Dette var temaer med høy aktualitet på denne tiden. Neyman-Pearson-teorien fra 1930-tallet hadde satt statistisk inferens inn i en beslutningsteoretisk ramme. Begrepsapparatet var videreutviklet og gitt et målteoretisk fundament på 1940 og 1950-tallet, i første rekke gjennom arbeider av Lehmann og Scheffé i Berkeley. Sverdrups hadde selv gitt bidrag til teorien. Han introduserte begrepet kompletthet og viste at en suffisient og komplett estimator er optimal. Den versjon av NP-teorien som han presenterte på konferansen kan karakteriseres med stikkordene generell, avrundet og elegant (se Sverdrup, 1967). Det var likevel rom for innvendinger fra de som var ubekvemme med rammeverket, f.eks. sto nærmere Fisher (se Fisher, 1955). Sverdrups forelesninger viste seg å være et



Arnljot Høyland (1924-2002)



Ole Barndorff-Nielsen (1935-)



Erling Sverdrup (1917-1994)

særdeles godt grunnlag for de mange med hovedfag i statistikk fra Blindern som tok fatt på doktorgradsstudier i Berkeley i årene etter.

Tidlig 1960-tallet var også tiden for den gryende neo-Bayesianisme, der Bruno de Finetti (1906-1985) og Leonard Savage (1917-1971) i årene før hadde utredet et Bayesiansk paradigme for statistisk inferens ut fra aksiomer for rasjonell adferd, og der Howard Raiffa og Robert Schlaifer hadde utviklet de mer tekniske sidene ved teorien. Gustav Elfving påtok seg oppgaven å gi et innblikk i denne utviklingen. Han var internasjonalt orientert, og var blitt god venn med Savage som hadde stimulert hans interesse. I tillegg til oversikt over det som hadde skjedd de siste 15 årene, ga Elfving tilkjenne kritiske merknader og alternativer. Spesielt pekte han på ønsket robusthet, dvs. ved alternative spesifiseringer av apriori-fordeling. Man lyttet med interesse, men tiden var ikke moden for utbredt bayesianisme. Hans meget lesverdige forelesninger ble senere publisert i Skandinavisk aktuarietidsskrift (Elfving, 1968).

Det tredje temaet på konferansen, Item-analyse, dreide seg om statistisk teori for konstruksjon av tester, et felt med opprinnelse i psyometri. Problemet er grovt sagt å kunne redusere antall spørsmål i et spørreskjema, slik at de gjenværende kan anses nødvendig og tilstrekkelig for formålet. Denne problemstilling hadde potensiale for anvendelse på mange områder, herunder utdanning, helse og markedsforskning. Georg Rasch var anerkjent som en pioner på dette feltet (se f.eks. Rasch, 1961), og de klassiske statistiske modellene hadde fått navn etter han. Det var derfor naturlig å la han presentere status i teorien for item-analyse, som på mange måter var blitt en dansk paradegren. Dette var før teorien for generaliserte lineære modeller var utviklet, der modellene nå kan sees på som spesialtilfeller.

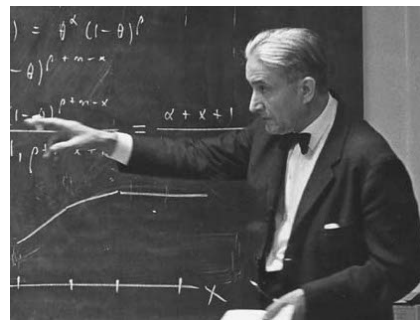
Den første konferansen hadde 70 deltakere (26 danske, 10 norske, 29 svenske og 5 finske). Ved avslutningen av konferansen ble det vedtatt ved akklamasjon å etablere den nordiske konferansen i matematisk statistikk som fast tradisjon, med sikte på å styrke matematisk statistikk som vitenskap i de nordiske land. Det neste møtet fant sted allerede året etter i 1966 i Lofthus i Hardanger.

Litteratur

Elfving, G (1968) Bayes statistics Part I-II, Skandinavisk Aktuarietidsskrift, 51, 134-155.

Fisher, R. A. (1955) Statistical method and scientific induction, Journal of the Royal Statistical Society, Ser. B, vol. 17, 69-78.

Rasch, G. (1961) On general laws and the meaning of measurement in psychology, Symposium on Mathematical Statistics and Probability IV, Berkeley CA: University of



Gustav Elfving (1908-1984)



Georg Rasch (1901-1980)

California Press.

Sverdrup, E. (1967) The present state of the decision theory and Neyman-Pearson theory, *Review of the International Statistical Institute* 34, 309–333.

Kontingent

Sondre Hølleland | Bergen

Medlemskontingenten er **kr 250** per år. Studentmedlemmer og livstidsmedlemmer betaler ikke kontingent. Kontingent for inneværende år, og eventuelt utestående kontingent for tidligere år, bes betales innen **30.april 2020**. Merk at kontingenten til og med 2017 var kr 200 per år.

Kontingenten kan betales på to måter:

- Ved innbetaling til kontonummer 0530.27.73299.
- Ved bruk av Vipps, til Vipps-nummer 97112.

I begge tilfeller ber vi om at du registrerer hvem og hvilket/hvilke år betalingen gjelder for.

Ta kontakt med kasserer Sondre Hølleland (sondre.holleland@uib.no) ved spørsmål angående kontingent.

Nytt fra NTNU

Nye mastergrader

- Karine Hagesæther Foss: Spatio-Temporal Gaussian Processes and Excursion Sets for Adaptive Environmental Sensing using Underwater Robotics. Veileder Jo Eidsvik.
- Øyvind Klåpbakken: Map Matching using Hidden Markov Models. Veileder Jo Eidsvik.
- Håkon Meyer: An Application of Statistical Learning in Direct Marketing Response Modelling. Veileder John Tyssedal.

Nye PhD-grader

- Yihan Cao
'Statistical methods for estimating fluctuating selection'. Disputas 3 April, 2020. I komiteen var Jarrod Hadfield, Univ of Ediburgh, UK, Hans Julius Skaug, Univ i Bergen, Ingelin Steinsland, NTNU (intern). Prøveforelesningstema: 'Hamiltonian Monte Carlo'

Nytilsatte PhD studenter

- Umut Altay. Veileder Geir-Arne Fuglstad

Nytt fra UiB

Nyansatte

- Konstantinos Giannakis, Post.doc. i ERC-prosjektet til Iain Johnston.

Nytt fra Matematisk institutt ved Universitetet i Oslo

Nye mastergrader

- Kim André Arntsen A functional approach to forward rate modelling Veileder: Fred Espen Benth
- Joel Fomete Tankmo Comparison of cohort and period life expectancy between countries in Scandinavia and the Mediterranean Veileder: Ørnulf Borgan

Nye PhD-grader

- Andreas Brandsæter
Data-driven methods for multiple sensor streams, with applications in the maritime industry
Veiledere: Ingrid K. Glad, Erik Vanem, Geir O. Storvik, Magne Aldrin, Arne B. Huseby.
Bedømmelseskomité: Associate professor Piero Baraldi, Politecnico di Milano, Associate Professor Prasad Lokukaluge Perera, UiT Norges arktiske universitet, Associate Professor Riccardo De Bin, Universitetet i Oslo
Digital disputas: 26. mars 2020

Nyansatte

- Ved Seksjon 2 (statistikk og data science) er Johan Pensar ansatt som førsteamanuensis, Haakon C. Bakka som postdoc og Marie Lilleborge som førstelektor.
- Videre er Thomas Minotto ansatt som stipendiat ved Seksjon 2 og Mari Dahl Eggen er ansatt som stipendiater ved Seksjon 3 (risiko og stokastikk).