

# Från en unik statistisk världskongress i Uppsala

M. Ribe



R&D Report  
Statistics Sweden  
Research - Methods - Development  
1990:13

## INLEDNING

### TILL

**R & D report : research, methods, development / Statistics Sweden. – Stockholm : Statistiska centralbyrån, 1988-2004. – Nr. 1988:1-2004:2.**

**Häri ingår Abstracts : sammanfattningar av metodrapporter från SCB med egen numrering.**

#### **Föregångare:**

Metodinformation : preliminär rapport från Statistiska centralbyrån. – Stockholm : Statistiska centralbyrån. – 1984-1986. – Nr 1984:1-1986:8.

U/ADB / Statistics Sweden. – Stockholm : Statistiska centralbyrån, 1986-1987. – Nr E24-E26

R & D report : research, methods, development, U/STM / Statistics Sweden. – Stockholm : Statistiska centralbyrån, 1987. – Nr 29-41.

#### **Efterföljare:**

Research and development : methodology reports from Statistics Sweden. – Stockholm : Statistiska centralbyrån. – 2006-. – Nr 2006:1-.

R & D Report 1990:13. Från en unik statistisk världskongress i Uppsala / Martin Ribe.  
Digitaliserad av Statistiska centralbyrån (SCB) 2016.

# Från en unik statistisk världskongress i Uppsala

M. Ribe



R&D Report  
Statistics Sweden  
Research - Methods - Development  
1990:13

Från trycket

Producent

Ansvarig utgivare

Förfrågningar

Oktober 1990

Statistiska centralbyrån, utvecklingsavdelningen

Åke Lönnqvist

Martin Ribe, tel. 08-783 48 54

© 1990, Statistiska centralbyrån

ISSN 0283-8680

Printed in Sweden

## *Från en unik statistisk världskongress i Uppsala*

M. Ribe

Under en vecka i mitten av augusti 1990 hölls en stor matematisk-statistisk konferens i Uppsala. Den var arrangerad av de två ledande sammanslutningarna i världen på området, nämligen dels Bernoulli Society, dels Institute of Mathematical Statistics (IMS). Detta var blott andra gången som Bernoulli Society ordnade en stor världskongress för allt inom matematisk statistik; den första var i Tasjkent för fyra år sedan. Uppsala-konferensen var också allra första gången som IMS höll sin ordinarie årskonferens tillsammans med Bernoulli-kongressen. Annars brukar IMS alltid ha sitt årsmöte tillsammans med American Statistical Association (ASA) i någon nordamerikansk storstad.

Det var alltså verkligen en unik manifestation som ägde rum i Uppsala. Det bjöds på en enorm mängd högintressanta föredrag inom olika fält. Inte mindre än ca 500 föredrag hölls, en imponerande hög aktivitet. Så gott som hela tiden pågick åtta parallella sessioner - och det under en hel vecka, från måndag till lördag!

Även kvalitativt var utbudet i absolut toppklass, med både allmänna översikter och färska forskningsresultat inom speciella områden. Utöver de temabundna sessionerna gavs ett antal fristående speciella föreläsningar, hållna av särskilt prominenta inbjudna talare. Sessionerna var annars på vanligt sätt inriktade på olika teman, var och en med ett antal föredrag, dels längre inbjudna (invited), dels korta av självanmälda talare (contributed). Särskilt specialföreläsningarna och många av de inbjudna föredragen gav utmärkta översikter över sådant som är på gång.

Konferensschemat verkade vara gjort med stor omsorg, med väl sammanpuslade sessioner och utan döda punkter. Även i övrigt flöt arrangemangen bra, med huvuddelen av konferensen i universitetets fräscha humanistisk-samhällsvetenskapliga centrum, och med inledande och avslutande plenarsessioner i den pampiga universitetsaulan. Bankett ingick även, enligt god Uppsala-sed förlagd till slottet.

Denna rapport kan naturligtvis bara spegla några personliga intryck av ett ringa urval av vad som bjöds på konferensen. Viss, men inte uteslutande, prioritet har givits åt sådant som kan ha potentiellt intresse med tanke SCB-statistik och individstatistik. Jag har även försökt ta med viss översiktlig bakgrundsinformation. Den som önskar mera detaljerade bakgrundsupplysningar om begrepp och annat hänvisas till det utomordentliga uppslagsverket

Encyclopedia of the Statistical Sciences. - Bengt Rosén och Eva Elvers har gjort värdefulla påpekanden som jag är tacksam för.

### **Mycket tillämpningar - men mest universitetsfolk**

Det hela var i högsta grad internationellt. Av de ca 750 deltagarna kom något hundratal från Sverige, ungefär lika många från USA, och resten från alla delar av världen. Notabelt var inslaget från Sovjet, efter de nya möjligheterna till utresa. Programmet tydde vidare på att Danmark är ett land där påfallande mycket händer inom den matematiska statistiken.

Från SCB var vi endast tre som deltog - låt vara att det därutöver faktiskt var fler än så som en gång varit SCB-are. Sanningen att säga fanns det dock även i övrigt bara mycket få deltagare som inte hade sitt hemvist i universitetsvärlden. Det var nog ganska synd att statistiker på företag och myndigheter, i Sverige och annorstädes, inte mera tog tillfället i akt.

Bland ämnena fanns åtskilligt som var inriktat på specifika praktiska tillämpningar, som medicin, försöksplanering, kvalitetskontroll, försäkring, miljöstatistik, bildbehandling, demografi och genetik. Även bland de mera generella ämnena fanns nästan genomgående en inriktning mot praktiskt tillämpbara metoder, för t.ex. regressionsanalys, tidsserieanalys, överlevnadsanalys och analys av rumslig variation.

Man kan väl även konstatera att innehållet bar tämligen klart vittnesbörd om särarten av "matematisk statistik" jämfört med "statistik" i övrigt. Karakteristiskt för matematisk statistik är att matematiska slumpmodeller står i centrum. Det väsentliga inom området är att utveckla modeller med lämpliga egenskaper, som ger en grund för att statistiskt beskriva och analysera olika företeelser i verkligheten. Rena tillämpningar av befintliga analysmetoder är däremot mera perifera från den matematisk-statistiska forskningens utgångspunkt.

Som det ser ut idag finner man vidare att stickprovsteori, och i än högre grad surveyteori i övrigt, endast har en undanskymd plats i den matematiska statistiken. Denna avgränsning är knappast helt logisk utan får nog närmast ses som utslag av en tradition som kommit att bli rådande - även om tendensen synbarligen är internationell. I sak borde nog modern surveymetodologi mycket väl kunna vara ett viktigt "matematisk-statistiskt" område.

### **Modellvalskriterier**

En session var särskilt inriktad på kriterier för val av modell vid analys av data.

I data-analys arbetar man ofta på det sättet att man utgår från ett observerat datamaterial och försöker finna mönster i det. Ett typexempel är när man försöker finna en regressionsmodell som kan visa hur en viss variabel beror av vissa andra variabler, s.k. förklarande variabler. För varje förklarande variabel kan man skatta ett parametervärde som anger hur starkt denna variabel är relaterad till den beroende variabeln. Om man från början inte vet vilka förklarande variabler man bör ta med, eller vilket slags modell som är bäst, så försöker man ibland pröva sig fram till en modell som ger så god "anpassning" som möjligt till data. Men det är inte meningsfullt att utan vidare göra så. Genom att ta till en modell med tillräckligt många parametrar, alltså tillräckligt många förklarande variabler, kan man teoretiskt sett alltid uppnå en perfekt anpassning. Men den goda anpassningen kan vara rent artificiell, och modellen och parameterskattningarna kanske inte alls säger något om verkligheten.

Det här är välkänt, och därför har man sedan ett antal år utvecklat mera rättvisande kriterier för modellval. Idén är att ge mått på hur "bra" modellen är för det givna materialet, med hänsyn både till modellens anpassning och till antalet parametrar. Modellerna blir "straffade" efter antalet parametrar, så att modeller med alltför många parametrar inte har något för sin förhållandevis goda anpassning. Numera klassiskt är Akaike-kriteriet, eller AIC (Akaike Information Criterion), som bygger på just denna idé.

Sessionen om modellvalsmetoder handlade mest om olika alternativ till AIC. Ett sådant gavs i ett föredrag av Jorma Rissanen och bygger på att modellen skall väljas så att modellspecifikationen och parameterskattningarna kan återges med totalt så få informationsbitar som möjligt. Med exempel avseende modeller för utjämning av kurvor ville han grafiskt åskådliggöra hur denna princip leder till naturliga resultat.

Som han och andra av talarna framhöll är AIC inte konsistent. Med det menas här att risken för att man skall få en överparametriserad modell inte kan pressas under en viss nivå - oavsett hur stort materialet är. Det vill säga, om ett aldrig så stort material är genererat av en okänd, "sann" modell, och om man gör en analys av materialet för att komma underfund med modellen, så finns en risk att AIC ger en alltför stor modell. Den modell som analysen leder till kommer då att innehålla parametrar som egentligen saknar mening. Att en metod även på mycket stora material kan ge stor risk för stora fel är inte så bra.

Benedikt Pötscher gav en allmän översikt som jämförde och visade analogin mellan olika alternativ. Ett ganska känt alternativ till AIC är BIC, eller Bayesian Information Criterion. För stora modeller ger BIC jämfört med AIC en hårdare "bestraffning" för antalet parametrar. Risken för överparametrisering blir då mindre, och man har visat att BIC är konsistent under diverse specifika förutsättningar. - SCB-bekante David Findley från U.S. Bureau of the Census presenterade ett resultat för modellval vid tidsserieanalys.

## Felanvända modeller

Även om man i matematisk statistik intresserar sig mera primärt för själva metoderna än för deras olika användningar, så fanns en session om bruk och missbruk av statistiska modeller. Förste talare var David Freedman, som gick igenom några valda exempel på hur man i praktiska tillämpningar drar slutsatser med hjälp av statistik. Ett föredöme än idag är J. Snows studier vid mitten av 1800-talet i London, rörande kolerans spridning. Genom en noggrann analys av data kunde Snow övertygande visa att sjukdomen spreds via dricksvattnet och inte via luften, vilket man dittills trott.

Andra tillämpningar är mindre föredömliga. I t.ex. samhällsvetenskaperna är det numera vanligt förekommande att man nyttjar olika statistiska modeller, såsom regressionsmodeller och liknande, ibland av ganska sofistikerat slag. Dessa modeller bygger på specifika förutsättningar om oberoende m.m., men i tillämpningslitteraturen diskuteras i regel knappast alls om dessa förutsättningar verkligen är uppfyllda. Valet av variabler i modellen kan ägnas viss uppmärksamhet, däremot inte om modellens inre logik verkligen passar för sakproblemet.

De exempel som gavs var "inte de värsta" men innehöll klart tvivelaktiga punkter, och tydligen hade de ändå rönt respekt och uppskattning inom sina områden. Man har t.ex. velat analysera hur politiska åsikter i vissa frågor beror på vissa bakgrundsvariabler. Man har då gjort regressionsanalyser och signifikanstest med USAs 50 delstater som observationer - behandlade som om de vore oberoende.

I samma session talade Paul Holland om användning av prov på elever, för antagning till utbildning och liknande. Här kan finnas intrikata mättekniska frågor som det behövs närmare analyser för att avslöja. Det kan till exempel visa sig att fel svar på en viss fråga är starkt positivt korrelerat till bra resultat på provet som helhet! Att ha med en sådan fråga förefaller onekligen tvivelaktigt.

## Expertsystem, nätverk och superdatorer

För något år sedan talades det mycket om statistiska expertsystem på dator. Ett statistiskt expertsystem, eller kunskapsbaserat system (KBS), är en programvara som hjälper användaren att välja analysmetod, allt efter data och förutsättningar. Nu har väl den första entusiasmen svalnat något för dessa "automatiserade statistiker", men vissa spår av ambitionerna finns fortfarande.

John Nelder berättade ett sådant kunskapsbaserat system som byggs upp kring välkända statistiska analysprogramvaror, bland dem GLIM. Systemet skall kunna ge användaren råd vid valet av statistisk analysmetod. Han betonade att systemet endast är rådgivande. Användaren har hela tiden själv full insyn



i och kontroll över hur analyserna skall utföras. Det finns också möjlighet att gå tillbaka om man har glömt vad man höll på med.

Men datorutvecklingen kommer att betyda mycket annat för statistikens utveckling. Paul Tukey (ej att förväxla med John Tukey) och Lorraine Denby från Bell visade spektakulära exempel på vilka möjligheter som kan öppnas genom kraftfulla arbetsstationer, sammankopplade i nätverk och med stor grafisk kapacitet. De "vanliga" persondatorerna är nu på väg att få så stor kapacitet att gränsen mot arbetsstationer utsuddas. En vision målades upp där persondatorn förmedlar både elektronisk post och TV-telefonsamtal i "fönster" på skärmen. Helt obehindrat framställer och överför man exempelvis rörliga grafiska framställningar av data ur elektroniskt lagrad litteratur.

Redan idag ger avancerad program- och maskinvara möjligheter att arbeta med interaktiv, rörlig och tredimensionell grafik som kan ge en helt ny åskådlighet. Lorraine Denby visade ett exempel på analys av ett fysiologiskt förlopp, hur blodtrycket varierar under pulsslagen. Genom att med musen vrida och vända på diagrammen, göra valda snitt ur tredimensionella framställningar m.m. kan man få en ny klarhet om det studerade fenomenet.

Det hela var huvudsakligen gjort i Bell-laboratoriernas egna statistikpaket S på en kraftfull personlig arbetsstation. Större beräkningskapacitet kan behövas för exempelvis analysberäkningar till stöd för tolkning av tredimensionella diagram. Då kan det bli fråga om stora matriskalkyler, såsom s.k. singularvärdesuppdelning, som kan ta några minuter på en Cray superdator. Så nog kan superdatorer komma väl till pass för statistisk analys.

## **Grafteori och orsaksmönster**

Av betydelse för bland annat expertsystem är en mera grundläggande teori för grafer över hur olika faktorer påverkar varandra. En sådan graf består helt enkelt av ett antal punkter och sammanbindande pilar eller linjer. Punkterna representerar variabler, och pilarna (eller linjerna) anger hur variablerna påverkar varandra. Idén är välbekant för dem som känner till programmet LISREL, där sådana samband kan modelleras.

Från den teoretiska sidan har detta slags grafer studerats inte minst av danska statistiker. M. Frydenberg definierade s.k. kedjegrafer, som är en vid klass av sådana grafer. Han införde olika begrepp för beroende och oberoende mellan grupper av variabler i graferna, egenskaper av s.k. Markov-typ, och visade hur de är relaterade till varandra. - Steffen Lauritzen presenterade en specifik inferensmodell på grafer, med tillämpning på expertsystem, dock inte ett statistiskt sådant.

Judea Pearl ställde sig ett annat problem, nämligen hur man ur observationer på givna variabler kan dra slutsatser om orsaksföljden mellan variablerna. Han började med att polemisera lite mot statistikernas hårda, men i praktiken

kanske lite dubbelmoraliska, dogm att statistiska samband inte får tolkas som orsakssamband. Sedan konstruerade han en modell där variablerna påverkas av varandra, och av oberoende yttre störningar, enligt en graf med pilar mellan variablerna. Utifrån ett tillräckligt stort observationsmaterial, simulerat från denna modell, är det möjligt att "återskapa" grafen, att se vilka variabler som påverkar varandra och hur.

Experimentet kan nog vara filosofiskt tankeväckande för LISREL-användare och andra. Men frågan är om det inte snarast visar svårigheterna, hur stort material och vilka förutsättningar som krävs, för att man ens teoretiskt skall kunna dra slutsatser om orsaker ur statistiska data.

### Modeller med störande parametrar

En liten serie med specialföreläsningar hölls av den berömde statistikern Sir David Cox. Han gav en översikt över en del aktuella resultat i inferensteori, där inte minst han själv varit produktiv med centrala bidrag. Föreläsningarna behandlade till stor del modeller med störande parametrar, s.k. nuisance parameters. Dessa störande parametrar svarar mot sådana effekter som man egentligen inte är intresserad av att skatta. Men de finns där ändå och kan påverka det man studerar. Man skulle dock vilja slippa ifrån dem, eftersom materialet kan vara otillräckligt för att man skall kunna skatta både de intressanta och de störande parametrarna med tillräcklig precision. Cox gick igenom olika sätt att hantera problemet, genom vissa korrigeringar och approximationer.

Ett exempel på en modell med störande parametrar är den modell för dödlighet som har sitt namn efter Cox, nämligen Cox' regressionsmodell. Dödlighetens beroende av åldern kan där följa en kurva av godtycklig form. Denna kurvas nivå beror av de hälsopåverkande faktorer som man vill studera, såsom bostadsort, arbetsmiljö m.m. Hur starkt varje sådan faktor inverkar på dödligheten anges av en parameter som man vill skatta i modellen. Dessa parametrar är alltså de effekter som man är intresserad av. Däremot är man i detta sammanhang inte intresserad av formen på kurvan som anger åldersberoendet. Kurvformen svarar alltså mot de störande parametrarna.

Cox-modellen är även exempel på en s.k. halvparametrisk modell. En sådan kännetecknas av att den modellerar fenomen som styrs av dels ett begränsat antal parametrar, dels kurvor eller ytor av mer eller mindre godtycklig form. Normalt är det de förra parametrarna man är intresserad av. Jon Wellner visade i sitt föredrag hur man kan härleda effektiva parameterskattningar i sådana halvparametriska modeller. Det var matematiskt elegant och föreföll kunna ha många tillämpningar, även om det möjligen beräkningsmässigt kan bli något komplicerat i enskilda fall.

## Medicinsk meta-analys

Ingram Olkin och Fred Mosteller presenterade idéerna bakom vad som kallas meta-analys av medicinska data. I studier av medicinska effekter, t.ex. resultat av olika behandlingsmetoder, arbetar man ofta med ganska små material. Det kan därför vara svårt att i de enskilda studierna få statistiskt säkerställda utslag. Osäkerhetsintervallen blir långa, och slutsatserna vanskliga.

Å andra sidan är det totalt sett många studier som utförs världen över. Det finns inte mindre än 23 000 medicinska tidskrifter som löpande publicerar olika försöksresultat. Det torde vara uppenbart att det även för en viss sjukdom och behandlingsmetod kan bli ganska många studier.

Frågan ger sig själv: Om man vill bedöma effekten av t.ex. en viss riskfaktor eller en viss behandlingsmetod, kan man då inte kombinera resultaten av flera studier? Då borde man kunna dra säkrare slutsatser än av varje studie för sig. Detta är just idén med meta-analysen. En sådan ansats är uppenbart inte utan vanskligheter. Olika studier är gjorda på olika sätt och olika väl dokumenterade. Att då bara "slå ihop" dem och utan vidare analysera dem som om de vore ett enda material verkar betänkligt. Men tillämpat med försiktighet och på rätt sätt måste det anses statistiskt fullt försvarligt att kombinera olika underlag i litteraturen, och det är också den vägen man måste gå för att på bästa sätt ta till vara den information som finns.

En fråga i sammanhanget är när man kan "sluta", när beläggen för t.ex. effekten av en ny behandlingsmetod har blivit tillräckligt starka. När metoden är ny finns endast få studier, och sedan kommer efter hand allt flera. Därmed blir grunden för att dra slutsatser om metoden allt starkare. Det gäller då att ha en korrekt "stoppregel", ett kriterium på när man kan dra en tillräckligt bestämd slutsats. En slumpmässig anhopning av skenbart gynnsamma rapporter får inte okontrollerat leda till felslut.

Ett allmänt problem vid bedömning av medicinska försöksresultat är att publiceringen kan vara selektiv. När man i ett försök inte lyckas belägga att t.ex. en viss behandling har effekt, så kan det anses ointressant, och då kanske studien inte blir rapporterad. Denna selektion gör det svårt att statistiskt korrekt bedöma värdet av publicerade resultat. Även om en publicerad studie tyder på att en behandling verkligen har effekt, så kan det möjligen tänkas att det också har gjorts ett antal opublicerade studier som inte visat någon effekt. Att tala om det "lyckade" resultatet som statistiskt säkerställt kan då vara bedrägligt.

Möjligen kunde man notera något delade meningar om hur stor problemet med selektion är. Larry Hedges såg det som allvarligt och presenterade en modell för hur man skulle kunna skatta hur selektionen slår. Idén bygger på att man utgår från den observerade fördelningen av de rapporterade effekternas styrka, och ett normalfördelningsantagande för de ursprungliga försöksresultaten. Man kan då göra en parametrisk skattning av hur selektionen beror

av försöksresultatet, och därmed hur stor den totala selektionen är. Synbarligen har modellantagandena stark inverkan på resultaten, men ansatsen borde ha intresse för att belysa problemet.

### Tidsserier m.m.

Vanliga metoder för analys av tidsserier bygger i regel på att serien är "stationär", d.v.s. att den hela tiden varierar på samma sätt kring samma konstanta nivå. Många serier man möter i verkligheten är i sig själva inte stationära, utan de kan visa en utveckling med långsiktiga förändringar. Ett vanligt trick är då att man bildar serien av förändringar, differenser, mellan på varandra följande tidpunkter i serien. I vissa fall kan serien av differenser då betraktas som stationär.

Sören Johansen talade i en specialföreläsning, till minne av Harald Cramér, om "kointegration" av tidsserier. Att bilda serien av förändringar är ett specialfall av kointegration. Mera allmänt kan kointegration innebära viktade differenser mellan olika serier. Om det är serier som på ett naturligt sätt är relaterade till varandra, såsom ränta och växelkurs, så bör detta visa sig i att en lämpligt viktad differens tenderar att bli stationär. Kriterier och metoder gavs på hur stationära mönster uppnås och hur lämpliga vikter skattas.

Temat stationaritet kom igen på ett mera teoretiskt sätt i en annan specialföreläsning, av Olav Kallenberg. Han gick där igenom vissa starkare villkor, som går ut på att fördelning och samvariation av en följd av slumpvariabler inte beror på ordningsföljden mellan variablerna. Ett flertal nyare resultat om hur dessa begrepp är relaterade och kan användas togs upp. Denna föreläsning låg på ett mera grundläggande plan av relativt djup teori och var knappast direkt tillämpningsinriktad.

En session hölls även om överlevnadsanalys, d.v.s. metoder för att följa och analysera i vilken omfattning t.ex. personer överlever till olika åldrar och efter hand dör. Per Kragh Andersen gav en översikt över de metoder som på senare tid utvecklats för att analysera överlevnaden i relativt små material av t.ex. patienter i en viss sjukdom. Där kombineras bl.a. Cox-modellen med modern teori för slumpmässiga förlopp, särskilt s.k. punktprocesser och martingaler.

Richard Gill presenterade en modell med flera tidsdimensioner. Möjligen var den inte främst avsedd för individers överlevnad utan snarare för tillförlitlighetsanalys av tekniska komponenter, också det ett viktigt tillämpningsområde för överlevnadsanalysen. - Niels Keiding talade om skattning av ålders- och kohortspecifik dödlighet efter kontinuerlig tid.

Nancy Heckman hade utvecklat en metod att testa om en regressionsfunktion är en monoton funktion av en annan. För detta införde hon ett slags rangkorrelation för kontinuerliga variabler.

## Dynamiska system och kontroll

Dynamiska system är beteckningen på ett för närvarande påtagligt aktivt område, både på ett matematiskt plan och för tillämpningar inom t.ex. ekonomi. Ya. Sinai, ett av de stora namnen på området, inledde hela konferensen med en översikt över vad som är på gång inom den matematiska teorin.

Denna teori är nära kopplad till den även i massmedier uppmärksammade kaos-teorin, om hur slump-liknande mönster kan uppkomma i ett icke slumpberoende förlopp. Datorsimuleringar har visat på intrikata sådana kaos-fenomen, som man först så småningom börjar förstå med matematisk teori. Detta fält upplevs idag som ett spännande teoriområde, där en hel del kan komma att hända framöver. Möjligen kan dessa arbeten med tiden komma att påverka vår syn på slump över huvud taget.

Ett likaså spektakulär utveckling genom datorsimuleringar är vad som kallas cellulära automater. Idén är enkel: på en spräcklig yta får färgerna "kriga" och erövra mark från varandra enligt vissa spelregler. David Griffeth visade i ett extrainsatt inslag en serie diabilder från datorsimuleringar av sådana förlopp. Det visar sig att det kan uppkomma de mest oväntade mönster, som sedan växer och inte går bort.

Från den ekonomiska sidan gäller intresset inte minst hur man kan kontrollera och styra dynamiska förlopp. Härvid bygger man bland annat på teori för stokastiska differentialekvationer, som även är ett område av rent matematiskt intresse. Differentialekvationer är den matematiska formen för våra vanligaste naturlagar, och en differentialekvation som är stokastisk innehåller ett slumpberoende i en viss form.

Detta är användbart för modellering av ekonomiska förlopp. Sådana tillämpningar behandlades av bland andra Darrel Duffie, vars ämne gällde jämviktspriser på värdepappersmarknader. Han byggde där på en ny nyttofunktion, som avser att vara mera realistisk än de som vanligen används, genom att den inte behöver vara en summa över tiden.

## Mera framtidsvyer

Spännande perspektiv för statistiken som helhet öppnades i Wilfrid Kendalls specialföreläsning. Den handlade om hur dator kan användas för komplicerade algebraiska beräkningar på stokastiska differentialekvationer. Detta kan tillämpas på beräkningar rörande slumpförlopp och sannolikhetsfördelningar på mångdimensionella "ytor". I en förlängning skulle detta kunna ge möjlighet att beräkna fördelningar för nya komplicerade skattningar, som därmed kan bli praktiskt användbara.

Peter Whittle redogjorde i sin Neyman-föreläsning för hur statistiska modeller kan användas för modellering av neuron-nät. Det är fråga om modeller av hur man tänker sig att nervcellerna samverkar i hjärnan, och i framtiden kommer vi förmodligen att få se datorer konstruerade enligt samma princip. Vad man bland annat försöker modellera är hur en anordning av ett stort antal parallella men samverkande, enkla komponenter, neuronerna, kan tjäna till att känna igen och identifiera mönster, såsom ett ord eller bilden av ett föremål. Indata kan även innehålla störningar.

För att ge en teoretisk grund för sådana neuron-nät får statistiska metoder en naturlig roll. Framtida datorer enligt denna princip kan tänkas få betydelse för s.k. databasmaskiner och associativa minnen. Sådana maskiner kan i sin tur komma att ge nya möjligheter till snabb effektiv sökning och sammanfattning av bland annat statistisk information.

Konferensen avslutades med att David Kendall höll en föreläsning instiftad till minne av Andrey Nikolayevich Kolmogorov. Den handlade om just Kolmogorov själv. Namnet Kolmogorov är kanske allra främst förknippat med Kolmogorovs axiomsystem för sannolikheter, som publicerades 1933 och är en grundsten inom den matematiska statistiken. Föreläsningen belyste vilken mångsidig och nyskapande matematiker Kolmogorov var, med insatser inom vida fält, och han var även engagerad i skolundervisning.

Kanske mindre allmänt bekant är att Kolmogorov på senare år arbetade med att ge sannolikhetsläran en helt annan grundval, som bygger på komplexitet av sifferföljder. En lång följd av nollor och ettor kan i högre eller lägre grad ha eller sakna drag av regelbundenhet. Man kan mäta graden av "komplexitet", alternativt tolkad som "slumpmässighet" eller "oregelbundenhet", hos sifferföljden. Ett mått på komplexiteten är den storlek ett datorprogram skulle behöva ha för att kunna producera sifferföljden i fråga.

På så vis fås ett nytt slumpbegrepp. Att bevisa de matematiska slumplagarna i termer av detta nya slumpbegrepp blev för några år sedan ett aktivt forskningsområde. Ännu är det väl ovisst om detta betraktelsesätt kommer att leda till något mera fruktbart, och under alla förhållanden torde det behövas en betydande mognadstid för idéerna. Men föreläsningen blev på sätt och vis inte endast en nostalgisk återblick utan även en spännande utblick framåt.

## BERNOULLI SOCIETY

Bernoulli-sällskapet, eller fullständigare The Bernoulli Society for Mathematical Statistics and Probability, är en av fyra sektioner inom International Statistical Institute (ISI). Medlemskap i sektionen är öppet för alla intresserade. I medlemsavgiften ingår tidskriften International Statistical Review, informationsbulletinen ISI Newsletter och rabatt på vissa andra tidskrifter. Bernoulli Society håller möte i anslutning till ISI-sessionerna vartannat år. Förutom de kongresser, av vilka den i Uppsala i år var den andra, anordnar Bernoulli Society sedan längre tillbaka regelbundet dels European Meetings of Statisticians, dels Conference on Stochastic Processes. Adress för ansökan om medlemskap m.m. är

The Executive Secretary of the Bernoulli Society  
International Statistical Institute  
P.O. Box 950  
NL-2270 AZ Voorburg  
Nederländerna

**R & D Reports** är en för U/ADB och U/STM gemensam publikationsserie, som fr o m 1988-01-01 ersätter de tidigare "gula" och "gröna" serierna. I serien ingår även **Abstracts** (sammanfattning av metodrapporter från SCB).

**R & D Reports Statistics Sweden** are published by the Department of Research & Development within Statistics Sweden. Reports dealing with statistical methods have green (grön) covers. Reports dealing with EDP methods have yellow (gul) covers. In addition, abstracts are published three times a year (light brown /beige/ covers).

Reports published earlier during 1990:

- 1990:1      Calibration Estimators and Generalized Raking Techniques in Survey  
(grön)      Sampling (**Jean-Claude Deville, Carl-Erik Särndal**)
- 1990:2      Sampling, Nonresponse and Measurement Issues in the 1984/85  
(grön)      Swedish Time Budget Survey (**Ingrid Lyberg**)
- 1990:3      Om justering för undertäckning vid undersökningar med urval i "rum  
(grön)      och tid" (**Bengt Rosén**)
- 1990:4      Data Processing at the Central Statistical Office - Lessons from recent  
(gul)      history (**M. Jambwa**)
- 1990:5      Abstracts I - sammanfattning av metodrapporter från SCB  
(beige)
- 1990:6      Sequential Poisson Sampling from a Business Register and its  
(grön)      Application to the Swedish Consumer Price Index (**Esbjörn Ohlsson**)
- 1990:7      Kvalitetsrapporten 1990 - huvudrapport (NN)  
(grön)
- 1990:8      Kvalitetsrapporten 1990 - bilagor (NN)  
(grön)
- 1990:9      Datorstöd vid datainsamlingen (**Carina Persson, Marit Jorsäter**)  
(grön)
- 1990:10    Abstracts II - sammanfattning av metodrapporter från SCB  
(beige)
- 1990:11    Rapport från DATI-projektets produktionsprov i arbetskraftsunder-  
(grön)      sökningarna augusti 1989 - januari 1990 (NN)
- 1990:12    Conceptual modelling and related methods and tools for computer-  
(gul)      aided design of information systems (**Bo Sundgren**)

Kvarvarande **beige** och **gröna** exemplar av ovanstående promemorior kan rekvireras från Inga-Lill Pettersson, U/LEDN, SCB, 115 81 STOCKHOLM, eller per telefon 08-783 49 56.

Kvarvarande **gula** exemplar kan rekvireras från Ingvar Andersson, U/ADB, SCB, 115 81 STOCKHOLM, eller per telefon 08-783 41 47.