

Met het Kalman filter vooruit

S. J. Koopman

maart 2000

*Met het Kalman filter vooruit
Leert van uw fouten uit het verleden
Houdt de toekomst in het vizier
Dan bent u de "wijze" in het heden*

U zou nu een verhandeling kunnen verwachten over de rol die het Kalman filter kan spelen in het maatschappelijk leven. Ik zou hier, vanmiddag, de stelling kunnen verdedigen dat als iedereen zijn of haar beslissing, op welk gebied dan ook, neemt op basis van het Kalman filter, het leven er stabiel uit zal zien. Het Kalman filter zegt ons immers dat we moeten leren van onze fouten uit het verleden. Het ter harte nemen van een gebeurtenis uit het verleden, waarop wij niet eerder hadden kunnen anticeren. Ons weerbaar maken tegen zo'n gebeurtenis, voor het geval dat het ons nogmaals zou overkomen. Dit zijn wezenlijke onderdelen van het leven. Zeker, deze levenshouding zal binnen behoudende maatschappelijke stromingen een gewillig oor vinden en dit kan mede verklaren waarom ik u, hier, vanuit deze positie toespreek. Maar, toegegeven, je leven laten leiden door het Kalman filter is minder spannend want de kleur zit vaak in de onvoorspelbaarheid. En het is juist de onvoorspelbaarheid die het Kalman filter tracht te minimaliseren. Hoe het ook zij, het leerproces is een belangrijk onderdeel van ons leven, zeker binnen deze muren. Voordat ik mij verga aan dwaze stellingen als dat het Kalman filter een leidraad is voor een beter leven, lijkt het mij beter om ons te richten op de verfrissende wiskundige werkelijkheid.

De vergelijkingen van het Kalman filter kunnen worden weergegeven als

$$\begin{aligned}a_{t+1} &= T_t a_t + T_t P_t Z_t' (Z_t P_t Z_t' + H_t)^{-1} (y_t - Z_t a_t), \\ P_{t+1} &= T_t P_t T_t' - T_t P_t Z_t' (Z_t P_t Z_t' + H_t)^{-1} Z_t P_t T_t' + R_t Q_t R_t',\end{aligned}$$

voor een bepaalde tijds-index t . Het symbool a_{t+1} representeert de toestand van morgen welke lineair kan worden afgeleid van de toestand van vandaag a_t plus de nieuwe informatie y_t die vandaag tot ons is gekomen. De onzekerheid van de toestand van morgen wordt gerepresenteerd door het symbool P_{t+1} welke kan worden afgeleid van de onzekerheid van vandaag P_t en de onzekerheid Q_t die inherent aan het leven zelf is. De symbolen T_t , Z_t , R_t en H_t representeren de lineaire ordening van de informatie die tot ons komt, de dynamiek van het leven en de interactie daartussen.

Hoewel ik er voor zou kunnen kiezen om allerlei maatschappelijke, politieke en psychologische beslissingsprocessen te beschrijven met behulp van het Kalman filter, lijkt het mij verstandiger en, gezien mijn wetenschappelijke contributies tot dusver, toepasselijker om een meer inhoudelijke beschouwing te geven over de rol van het Kalman filter in de statistische analyse van waarnemingen die over tijd gemeten worden (tijdreeksen) binnen een economische context.

Inleiding

Aan het begin van deze lezing zal ik twee verschillende methodologiën voor tijdreeksanalyse kort introduceren waarna ik de problematiek in een meer historische context zal plaatsen. Vervolgens zal ik beide methodologiën met elkaar vergelijken en ik zal proberen aan te tonen dat de state space methodologie te prefereren is. Daarna zal ik het brede werkterrein van de state space met u verkennen met een speciale referentie naar de rol die het kan spelen in de econometrie. Ik hoop dat ik u en vooral de studenten onder u kan overtuigen dat er veel werk te doen is op dit terrein. Ik zal afsluiten met het uitspreken van een voorspelling.

Deze lezing gaat over hetgeen in het Engels bekend staat als *the state space methodology* en in het Nederlands soms wordt aangeduid als *de toestandsruimte methodologie*. Hoewel ik in deze rede de statistische begrippen zoveel mogelijk in het Nederlands zal aanduiden, zal ik de Engelse aanduiding "state space" handhaven omdat ik mij niet kan vinden in de Nederlandse vertaling hiervan hoe correct deze in principe ook is.

De traditionele decompositie van een tijdreeks is gegeven door

$$\begin{aligned} \text{observatie} &= \text{trend} + \text{seizoen} + \text{storing} \\ y_t &= \mu_t + \gamma_t + \varepsilon_t, \end{aligned}$$

Het onderliggende idee van state space tijdreeksanalyse is eerst het modelleren van de onafhankelijke componenten trend, seizoen en storing, om vervolgens de componenten bij elkaar te voegen in een lineair model. De algemene vorm van het model is

$$\begin{aligned} y_t &= Z_t \alpha_t + \varepsilon_t, & \varepsilon_t &\sim N(0, H_t), \\ \alpha_{t+1} &= T_t \alpha_t + R_t \eta_t, & \eta_t &\sim N(0, Q_t), \end{aligned}$$

voor $t = 1, \dots, n$. De waarneming y_t is lineair afhankelijk van de state vector α_t waarbij we aannemen dat de rij-vector Z_t bekend is. De storing ε_t wordt niet geobserveerd maar we nemen aan dat deze gegenereerd is door een normale verdeling. Als de state vector α_t constant zou zijn dan reduceert het model zich tot een gewoon lineair regressie model. Maar we zien dat de state vector over tijd kan veranderen via een multivariaat autoregressief proces waarbij we aannemen dat de matrices T_t en R_t bekend zijn. De variantie matrices H_t en Q_t hangen veelal af van een aantal onbekende coëfficiënten die we moeten schatten.

Het Kalman filter is gekoppeld aan het state space model en het berekent de schatter a_{t+1} voor de onbekende state vector α_{t+1} op basis van de waarnemingen $y_1, y_2, \dots, y_t$. De schatter a_{t+1} heeft optimale statistische eigenschappen en kan worden gebruikt voor het berekenen van voorspellingen voor toekomstige observaties zoals $\hat{y}_{t+1} = Z_{t+1} a_{t+1}$.

De elementen van de state vector representeren de verschillende componenten van het model zoals trend, seizoen en storing. In andere situaties, componenten voor cycli bewegingen, kalender effecten, verklarende variabelen en interventie effecten kunnen geïntroduceerd worden in het model. Een eenvoudig voorbeeld van een state space model is het lokaal nivo model welke gegeven is als

$$y_t = \mu_t + \varepsilon_t, \quad \mu_{t+1} = \mu_t + \eta_t,$$

waarbij het nivo van de tijdreeks wordt gerepresenteerd door μ_t die verandert over tijd: het nivo van morgen is gelijk aan het nivo van vandaag plus een storing η_t . De storing η_t geeft de verandering aan van de situatie van dag tot dag en de storing ε_t representeert de onzekerheid van het onderliggende proces dat de observaties genereert.

Hoewel state space modellen bekend zijn onder de specialisten, ze worden nog maar weinig gebruikt in de praktijk ¹. De meeste toepassingen van tijdreeksanalyse in de statistiek en econometrie van de afgelopen twintig jaar zijn gebaseerd op de zogenaamde ARIMA modellen van Box en Jenkins ². Het basis idee van Box en Jenkins is om trend en seizoen effecten te elimineren via het nemen van verschillen. Voorbeelden van het nemen van verschillen zijn

$$\Delta y_t = y_t - y_{t-1}, \quad \Delta^2 y_t = \Delta(\Delta y_t), \quad \Delta_s y_t = y_t - y_{t-s},$$

waarbij we aannemen dat er s "maanden" zijn in een "jaar". Het nemen van verschillen gaat door totdat trend en seizoen effecten zijn geëlimineerd. Vervolgens wordt aangenomen dat de lokale statistische eigenschappen van de tijdreeks, na het nemen van verschillen, constant is over tijd. Wanneer de verwachting en variantie van Δy_t constant is voor alle perioden $t = 1, \dots, n$ dan duiden we de tijdreeks Δy_t aan als een *stationaire* tijdreeks en y_t kan gemodelleerd worden als een *autoregressive integrated moving average* (ARIMA) model. Een eenvoudig voorbeeld van zo'n model is

$$\Delta y_t = \nu_t + \theta \nu_{t-1}, \quad \nu_t \sim N(0, \sigma^2),$$

waarbij de coëfficiënt θ onbekend is en geschat dient te worden.

Voordat deze rede zal verworden tot een inleidend college tijdreeksanalyse zal ik de twee methodologieën niet verder toelichten. Ik heb met deze korte inleiding willen laten zien dat beide methoden (state space en Box-Jenkins) ogenschijnlijk zeer verschillend van elkaar zijn. Echter, wanneer beide methoden in hun historische context worden geplaatst, blijken zij van origine niet zo verschillend te zijn. De methoden worden conceptueel wel verschillend wanneer we meer realistische modellen gaan analyseren. Onder deze omstandigheden heb ik een sterke voorkeur voor de state space methodologie.

¹Eén van de uitzonderingen hierop is het STSM systeem van de Groot, Koopman en Ooms (1999) dat op het Directie Binnenlands Geldwezen van het Ministerie van Financiën wordt gebruikt voor het voorspellen van de dagelijkse belastingsinkomsten en de maandcumulatie ervan.

²Zie G.E.P Box en G.M. Jenkins (1970) "Time Series Analysis: Forecasting and Control", San Francisco: Holden-Day.

Voorspellen via exponentieel gewogen gemiddelde

In de jaren vijftig zijn verschillende methoden ontwikkeld voor het voorspellen van tijdreeksen op basis van exponentionele gewichten. Ik begin de expositie met een eenvoudig voorbeeld. In het diagram (1) ziet u de jaarlijkse groei van ons nationaal inkomen (in percentages) zoals deze gerapporteerd is door het Centraal Planbureau in Den Haag. Wanneer we met behulp van eenvoudige middelen, zonder gebruik te maken van zakrekenmachines en computers en zonder gebruik te maken van econometrische modellen, de groei van het nationaal inkomen willen voorspellen voor het volgende jaar dan kunnen we een simpele lijn door de observaties tekenen die zo goed mogelijk het patroon van de reeks volgt. We zijn geïnteresseerd in het voorspellen van toekomstige waarden. Dus het tekenen van de lijn of de curve dient in ons geval alleen rekening te houden met observaties uit het verleden: het uiteinde van de potlood probeert een zo goed mogelijk pad te volgen door de tijdreeks door slechts rekening te houden met de informatie die gepasseerd is. In wezen is dit het principe van het Kalman filter: de optimale statistische schatter ³ van de toekomstige observatie voor periode $t + 1$ bepalen op basis van informatie uit het verleden en heden.

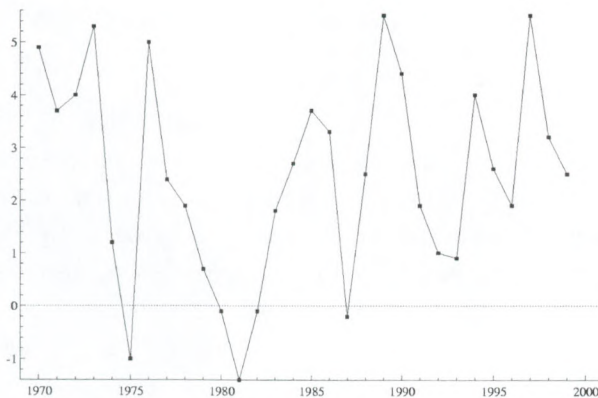


Figure 1: Groei nationaal Inkomen in Nederland

Bron: Centraal Planbureau te Den Haag

We nemen gemakshalve aan dat de beste voorspelling voor de volgende periode

³Optimaal wordt hier geïnterpreteerd als de minimale verwachte kwadratische voorspelfout.

gelijk is aan de waarde die de potlood aangeeft voor de huidige periode. Een mogelijkheid is om een lijn te tekenen door alle waarneempunten zodat de voorspelling voor de volgende periode gelijk is aan de huidige observatie. In veel gevallen lijkt dit niet verstandig want de situatie in de huidige periode zal niet precies gelijk zijn aan de situatie in de vorige periode ⁴. We willen echter wel rekening houden met ontwikkelingen die zich een aantal perioden aanhouden en daardoor meer structureel van karakter zijn. Men zal zo goed mogelijk een lijn tekenen die het patroon van de tijdreeks volgt.

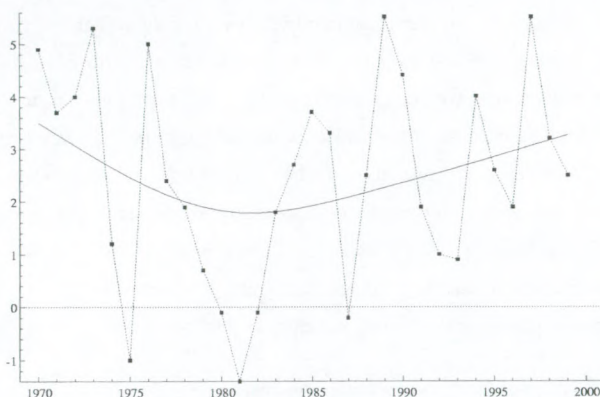


Figure 2: Een lijn door waarnemingen

In diagram (2) heb ik een poging gewaagd. Voorspellingen van de groei van het nationaal inkomen voor het volgende jaar en de daarop volgende jaren, zijn gegeven door het uiteinde van de lijn recht door te trekken (extrapoleren) naar toekomstige perioden. Met het tekenen van een dergelijke lijn probeert men, wellicht onbewust, het gemiddelde van het recente verleden zo goed mogelijk te benaderen. De keuze van wat we bedoelen met het recente verleden is van wezenlijk belang hier. Wanneer met een ver verleden wordt rekening gehouden, dan zal de lijn van periode tot periode niet zoveel veranderen dan wanneer met een kort verleden wordt rekening gehouden. Dit is natuurlijk slechts ten dele waar omdat de verandering van de lijn ook afhangt

⁴Een bekende uitzondering op deze regel is een tijdreeks van prijzen die op volledig doorzichtige en efficiënte markten, zoals deze bestaan voor aardappelen, bloemen en aandelen, tot stand zijn gekomen. Economische theorie suggereert dat alle relevante handelsinformatie in de prijs verdisconteert is. De huidige prijs is daardoor de beste voorspelling voor de prijs van morgen.

van de karakteristieken van de tijdreeks zelf. Deze twee aspecten zijn met elkaar verweven. Wanneer het patroon van de waarnemingen veel verandert over tijd dan is het raadzaam om geen rekening te houden met een ver verleden omdat die informatie dan niet meer relevant is voor de nabije toekomst. Wanneer een tijdreeks weinig verandert dan maakt het eigenlijk niet zoveel uit want de informatie uit een recent verleden of die uit een ver verleden is min of meer gelijk. Deze bespiegelingen zijn belangrijk maar vrijblijvend want of een reeks nu wel of niet veel verandert is uiteindelijk een subjectief gegeven. Het is het natuurlijke instinct van een econometrist om voor het voorspellen van toekomstige nog niet waargenomen observaties een objectief raamwerk te ontwikkelen.

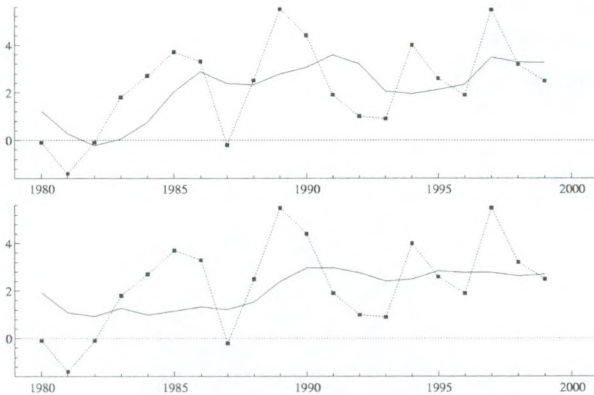


Figure 3: Voortschrijdend gemiddelden

Tot hoever het verleden van een tijdreeks nog relevant is voor de toekomst wordt bepaald aan de lengte van wat we zullen aanduiden als het venster. Het gemiddelde definiëren we als een rekenkundig gemiddelde. Het globaal gemiddelde is gelijk aan het rekenkundig gemiddelde van alle waarnemingen. Het voortschrijdend gemiddelde in een periode t wordt berekend als

$$\bar{y}_t = \frac{1}{p} \sum_{j=0}^{p-1} y_{t-j}, \quad t = p, \dots, n,$$

waarbij p de lengte van het venster is. De voorspelling voor de volgende periode is gelijk aan

$$\hat{y}_{t+j} = \bar{y}_t, \quad j = 1, 2, \dots,$$

en van speciaal belang is de voorspelling voor $t = n$ omdat dit de toekomst is waarvoor we geen observaties beschikbaar hebben. In diagram (3) zijn voortschrijdend gemiddelden weergegeven op basis van 5, 10 en alle observaties. Het is duidelijk te zien dat de verschillende curves anders zijn. We hebben nu een mechanische manier gevonden om een lijn door de tijdreeks te tekenen. Let wel, de lijn hangt af van de lengte van het venster en wordt alleen bepaald door observaties uit het verleden en heden.

De wijze waarop het voortschrijdend gemiddelde is berekend heeft als nadeel dat de invloed van waarnemingen schoksgewijs geïntroduceerd worden. Dit blijkt bijvoorbeeld wanneer een bepaalde waarneming y_j een zeer grote waarde heeft. De invloed van deze waarneming op het voortschrijdend gemiddelde blijft p perioden lang hetzelfde; buiten het venster om heeft deze waarneming geen enkele invloed. Het is alles of niets. Een minder schoksgewijze methode is het gewogen voortschrijdend gemiddelde: waarnemingen uit het verleden ontvangen minder gewicht wanneer deze verder gelegen zijn van de huidige periode. Dus in plaats van een rechthoekig gewichten patroon – diagram 4 (i) –, kunnen we ook een driehoeks gewichten patroon – diagram 4 (ii) – of een gewichten patroon gelijkend op een glijbaan – diagram 4 (iii) – gebruiken voor het berekenen van het gewogen gemiddelde ⁵.



Figure 4: Gewichten van voortschrijdend gemiddelden

⁵Dezelfde overwegingen bestaan voor niet-parametrische methoden in de econometrie en statistiek zoals Kernel methoden voor het schatten van kansverdelingen en het schatten van niet-lineaire verbanden tussen waargenomen variabelen. Over deze connecties zal ik later meer zeggen.

In formule vorm schrijven we

$$\bar{y}_t = \sum_{j=0}^{p-1} w_j y_{t-j}, \quad t = p, \dots, n,$$

waarbij

$$\sum_{j=0}^{p-1} w_j = 1,$$

zodat de gewichten het onderliggend nivo van de waarnemingen niet kunnen beïnvloeden. In het geval dat $w_j = 1/p$ dan is het gewogen rekenkundig gemiddelde gelijk aan het rekenkundig gemiddelde. Een bekende manier om gewichten te formeren is deze exponentieel te laten afnemen naarmate de afstand groter wordt ten opzichte van het huidige punt. In formule vorm schrijven we

$$w_j = (1 - \lambda)\lambda^j, \quad 0 < \lambda < 1,$$

waarvoor geldt dat

$$\sum_{j=0}^{p-1} w_j \approx 1.$$

De benadering wordt nauwkeuriger naarmate een hogere waarde voor p wordt gekozen. In wezen maakt de keuze voor p niet zo veel meer uit; de keuze is verplaatst naar een waarde voor de kortingsfaktor λ . Naarmate λ dichterbij 1 wordt gekozen, dient j een hogere waarde te hebben om w_j nul te laten naderen zodat deze verwaarloosbaar wordt; dit impliceert een hoge waarde voor p . Wanneer de niet-negatieve faktor λ klein wordt gekozen, dan daalt het gewichtenpatroon snel af tot nul hetgeen een kleine waarde voor p impliceert. Een aantal voorbeelden van gewichtenpatronen voor verschillende waarden van λ kunt u vinden in diagram 5.

Met het exponentieel gewichtenschema kunnen de een-stap vooruit voorspellingen worden berekend als

$$\hat{y}_{t+1} = \bar{y}_t = (1 - \lambda) \sum_{j=0}^{\infty} \lambda^j y_{t-j}.$$

waarbij we p hebben vervangen door het oneindige teken ∞ . Deze formulering is niet zo geschikt voor het daadwerkelijk uitrekenen van de voorspelling maar we kunnen gelukkigerwijs deze herschrijven als een recursie:

$$\hat{y}_{t+1} = (1 - \lambda)y_t + \lambda\hat{y}_t.$$

Deze formule is eenvoudig te berekenen via een zakrekenmachine: de nieuwe voorspelling hangt slechts af van de huidige waarneming y_t en de voorspelling van "gisteren"



Figure 5: Exponentionele gewichtenschema's voor $\lambda = 0.1, 0.25$ en 0.6

voor "vandaag", $\hat{y}_t$. Deze eenvoudige berekeningswijze, gekoppeld aan een goed gefundeerd verhaal, maakte deze methode zeer populair in de jaren vijftig en zestig in het bedrijfsleven. Er was veel lopende band werk op productie afdelingen van grote bedrijven. Het besluit om machines te laten stoppen voor onderhoud kwam veelal tot stand wanneer gedurende een aantal perioden de voorspellingen, zoals deze hierboven uitgerekend werden, te optimistisch waren. Dergelijke eenvoudige methoden werden ook gebruikt voor het voorspellen van verkoopcijfers van allerlei producten. De effectiviteit van een prijsactie werd bepaald door de gerealiseerde verkoopcijfers te vergelijken met de eerder berekende voorspellingen. Ik wil hier verder opmerken dat de exponentionele voorspelmethode nog steeds zijn opgang doet en soms uit een onverwachte hoek. Bijvoorbeeld, het exponentionele gewichtenschema wordt prominent gebruikt in het populaire programma RiskMetricsTM van de internationale bankinstelling J.P. Morgan voor het risico-management van financiële instellingen ⁶.

⁶Drs. R. van de Goorbergh (KUB) heeft in zijn afstudeerverslag over risico-management via Value-at-Risk analyses laten zien dat alleen kleine verbeteringen ten opzichte van RiskMetricsTM kunnen worden verkregen en dan alleen door gebruik te maken van computer-intensieve methoden.

Schatten van kortingsfaktor

De keuze van λ is belangrijk in de exponentionele gewichten methode. De voorspellingen hangen af van de waarde voor λ . Wanneer deze waarde te hoog is, dan wordt te veel rekening gehouden met het verleden en dat kan tot voorspellingen leiden die te traag reageren op nieuwe ontwikkelingen. Wanneer deze waarde te laag is, dan wordt juist te snel gereageerd op ontwikkelingen die wellicht eenmalig van karakter zijn geweest. Hoe kunnen we nu een optimale waarde voor λ vinden? In de beste traditie die we in de statistiek en econometrie hebben, kunnen we λ zo kiezen dat de som van de kwadratische voorspelfouten uit het verleden minimaal is:

$$\text{minimaliseer } \sum_{t=p^*}^n (y_t - \hat{y}_t)^2 \text{ met betrekking tot } \lambda,$$

voor een bepaalde waarde van p^* . Deze minimaliseringsopdracht is relatief eenvoudig uit te voeren door voor verschillende waarden van λ , de som van kwadratische fouten uit te rekenen en vervolgens die waarde van λ te kiezen waarvoor de som minimaal is. De keuze voor λ kan nauwkeurig bepaald worden door voor veel verschillende λ 's de som te berekenen of door slim te minimaliseren via geavanceerde numerieke zoekmethoden.

In ons voorbeeld van percentuele groei van het nationaal inkomen kwam ik tot het volgende resultaat: λ moet ongeveer gelijk zijn aan 0.25 want voor deze waarde was de kwadraten som gelijk aan 118 en voor andere waarden van λ was de som hoger. Bijvoorbeeld, voor de λ waarden 0.1 en 0.6 was de kwadraten som gelijk aan 121 en 126, respectievelijk.

Het statistische model

Een belangrijke contributie is gemaakt door Professor J.F. Muth in 1960 die liet zien dat met behulp van statistische analyses op basis van een model, voorspellingen op dezelfde wijze gegenereerd worden als via het exponentieel gewogen gemiddelde ⁷. Muth gebruikte het model

$$\begin{aligned} y_t &= \mu_t + \varepsilon_t, \\ \mu_{t+1} &= \mu_t + \eta_t, \end{aligned}$$

⁷Zie J.F. Muth (1960), Journal of the American Statistical Association 55, p.299-305.

waarbij de onafhankelijke storingstermen ε_t en η_t onderling onafhankelijk zijn en verwachting nul hebben met constante varianties σ_ε^2 en σ_η^2 , respectievelijk ⁸. Dit model is gelijk aan het meest eenvoudige state space model dat we eerder hebben aangeduid als het lokaal nivo model. In het boek dat ik heb geschreven met Professor James Durbin en dat later dit jaar zal verschijnen, wijden we het eerste hoofdstuk volledig aan dit eenvoudige model ⁹. We leiden via de bekende regels van iteratieve verwachtingen een expressie voor de geconditioneerde verwachting af voor het nivo μ_{t+1} gegeven de gerealiseerde waarnemingen y_j uit het verleden ($j = 1, \dots, t-1$) en de huidige waarneming ($j = t$). Het blijkt dat deze geconditioneerde verwachting, welke we zullen aanduiden als m_{t+1} , berekend kan worden via een recursie die wij kennen als het Kalman filter. Na enige tijd (zeg, $t \geq p^*$) convergeert het Kalman filter naar:

$$m_{t+1} = m_t + \bar{K}(y_t - m_t), \quad t = p^*, \dots, n,$$

waarbij $\bar{K}$ een functie is van

$$q = \frac{\sigma_\eta^2}{\sigma_\varepsilon^2},$$

hetgeen we zullen aanduiden als de *nivo tot storing ratio*. De voorspelling voor de waarneming van morgen is gelijk aan de conditionele verwachting m_{t+1} :

$$\hat{y}_{t+1} = m_{t+1}.$$

Het is eenvoudig door herschikking te laten zien dat het algoritme van het voortschrijdend exponentieel gewogen gemiddelde kan worden herschreven als

$$\hat{y}_{t+1} = \hat{y}_t + (1 - \lambda)(y_t - \hat{y}_t),$$

en dit is gelijk aan het geconvergeerde Kalman filter voor het lokaal nivo model met $\bar{K} = 1 - \lambda$ hetgeen impliceert dat er een verband bestaat tussen λ en q omdat $\bar{K}$ slechts een functie is van de nivo tot storing ratio q .

Ik zal het belang van het resultaat van Professor Muth kort samenvatten.

- Het voortschrijdend exponentieel gewogen gemiddelde levert ons een voorspelling op voor de volgende periode. In vele situaties is het wenselijk om ook

⁸Ik ga in deze rede voorbij aan het initialisatie probleem voor het proces van μ_t . Het model is pas volledig gedefinieerd wanneer een aanname is gemaakt over μ_1 . We kunnen μ_1 als een onbekende coëfficiënt beschouwen en deze schatten via een regressie procedure.

⁹Zie Durbin en Koopman (2000), "Time Series Analysis by State Space Methods", Oxford University Press.

een indicatie te geven over de waarschijnlijkheid van de voorspelling. Door kansverdelingsfuncties toe te kennen aan de storings van het lokaal nivo model en daarmee rekening te houden in de statistische analyses kunnen we uitspraken doen over de waarschijnlijkheid van de voorspelling.

- De onbekende coëfficiënt q kan op statistische gronden bepaald worden ofwel de meest aannemelijke schatter van q kan berekend worden wanneer kansverdelingsfuncties zijn toegekend aan de storingstermen. Dit komt overeen met de keuze van een waarde voor de λ coëfficiënt van het exponentieel gewogen gemiddelde. Statistische eigenschappen van de meest aannemelijke schatter zijn bekend onder de conditie dat heel veel observaties worden waargenomen.
- Het model kan worden aangepast en uitgebreid op een systematische wijze. Ik zal hier later op terugkomen.
- Statistische toetsen kunnen worden gebruikt om te bepalen of het model adequaat is.

In de Box-Jenkins methode zal men het model van Muth niet in deze vorm kunnen gebruiken omdat het model impliceert dat de waarnemingen y_t niet stationair zijn: de ongeconditioneerde verwachting van y_t is gelijk aan μ_t en deze is niet constant over tijd. De Box-Jenkins methode tracht, door het nemen van verschillen, een reeks te creëren dat wel stationair is. In dit geval kan men zich beperken tot de groei van de tijdreeks welke gelijk is aan de eerste verschillen van de waarnemingen y_t :

$$\Delta y_t = y_t - y_{t-1}, \quad t = 2, \dots, n.$$

Als we aannemen dat de waarnemingen zijn gegenereerd door het lokaal nivo model, dan kan het model voor de groei worden weergegeven als

$$\Delta y_t = \varepsilon_t - \varepsilon_{t-1} + \eta_{t-1}.$$

Het vergt weinig moeite om te laten zien dat de ongeconditioneerde verwachting en variantie van de groei Δy_t wel constant zijn over tijd. We concluderen dat de tijdreeks stationair is na het nemen van eerste verschillen. Verder, alle auto-correlaties voor Δy_t zijn nul behalve de auto-correlatie van één periode vertraging is ongelijk aan nul. Dit is gelijk aan de auto-correlatie functie van Δy_t in het eenvoudige ARIMA model dat

we eerder schreven als

$$\Delta y_t = \nu_t - \theta \nu_{t-1}, \quad 1 > \theta > 0,$$

met onafhankelijke storingsterm ν_t . Dit model wordt wel aangeduid als het AR-IMA(0,1,1) model voor y_t . De coëfficiënt θ is een functie van q en dus weer een functie van λ . Dit leidt tot de conclusie dat het voortschrijdend exponentieel gewogen gemiddelde twee model representanten heeft: het lokaal nivo state space model en het Box-Jenkins ARIMA(0,1,1) model. Het feit dat twee verschillende modellen dezelfde voorspellingen kunnen genereren via het voortschrijdend exponentieel gewogen gemiddelde is een interessante observatie en kan van praktische waarde zijn.

Het gladstrijken van tijdreeksen

Voor de illustratie van de jaarlijkse groei van het nationaal inkomen in Nederland heb ik de meest aannemelijke schatter voor de nivo tot storing ratio q van het lokaal nivo model gevonden om vervolgens de voorspellingen m_{t+1} te berekenen voor $t = p^*, \dots, n$. De voorspellingen tezamen met een 50% betrouwbaarheidsinterval zijn gepresenteerd in diagram 6¹⁰. De lijn door de observaties geeft een weergave van de ontwikkelingen over tijd betreffende de groei van het nationaal inkomen. Deze ontwikkeling is duidelijker in beeld gebracht dan door de waarnemingen zelf omdat de lijn m_t geleidelijk verloopt over tijd: de schokken zijn geëlimineerd omdat m_t een lokaal gemiddelde is van observaties uit het verleden.

Een beter beeld van de ontwikkelingen over tijd kan worden verkregen door het gemiddelde te berekenen op basis van observaties uit het verleden, het heden en de toekomst, tot hoever deze beschikbaar zijn:

$$\tilde{y}_t = \sum_{j=-p}^p w_j y_{t+j}.$$

Dit wordt wel aangeduid als *smoothing* en we zullen dit vrij uit het Engels vertalen als *gladstrijken*. In de tijdreeks literatuur wordt het ook wel aangeduid als *signal extraction* of *trend extraction* hetgeen moeilijk vertaald kan worden.

In diagram 7 (ii) ziet u een voorbeeld van een exponentieel gewichtenpatroon dat zowel gewichten allocceert aan observaties uit het verleden als aan observaties in de

¹⁰Ik heb gekozen voor een 50% betrouwbaarheidsinterval omdat deze naar mijn mening een grotere praktische waarde heeft voor voorspellingen dan de gebruikelijke 95% betrouwbaarheidsinterval.

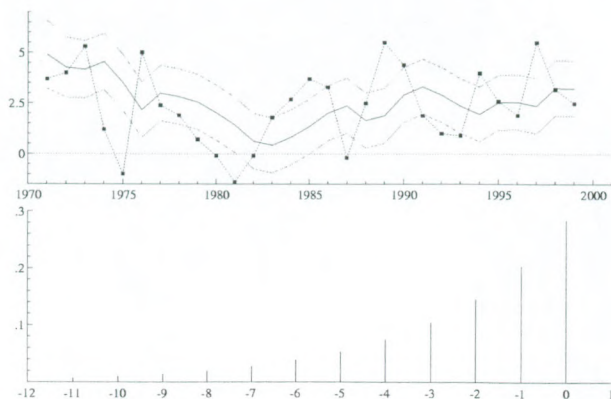


Figure 6: Voorspellingen met 50% betrouwbaarheidsinterval en gewichten ($q = 0.11$)

toekomst. Ik zal hier niet verder ingaan op hoe de gewichten berekend kunnen worden maar de exponentionele gewichten hangen wederom af van de coëfficiënt $0 < \lambda < 1$ ¹¹. Het lokaal nivo model

$$y_t = \mu_t + \varepsilon_t, \quad \mu_{t+1} = \mu_t + \eta_t,$$

kan nog steeds gebruikt worden als het onderliggende statistisch model. Echter, in plaats van alleen het Kalman filter te gebruiken voor het "voorwaarts" berekenen van de conditionele verwachting dienen we nu ook een "achterwaarts" algoritme te gebruiken om de conditionele verwachting voor μ_t gegeven alle waarnemingen y_j ($j = 1, \dots, n$) te berekenen. Diagram 7 (i) toont het onderliggende nivo van de groei van het nationaal inkomen met een 95% betrouwbaarheidsinterval en het is duidelijk te zien dat de Nederlandse economie in de afgelopen drie decennia een stabiele groeiperiode heeft gekend¹². Alleen aan het begin van de jaren tachtig van de vorige eeuw was er sprake van een lagere groei in de economie.

Het gladstrijken van tijdreeksen op basis van een statistisch model verdient mijn inziens de voorkeur zoals ik zojuist heb geargumenteed. Het is de taak van statistici en econometristen om redelijke modellen voor te dragen die door middel van

¹¹Voor een algoritme voor het berekenen van gewichten voor schatters van onbekenden in een state space model, zie S.J. Koopman en A.C. Harvey (1999), Discussion paper, Vrije Universiteit.

¹²Hier kies ik voor een 95% betrouwbaarheidsinterval omdat we hier niet met voorspellingen van doen hebben.

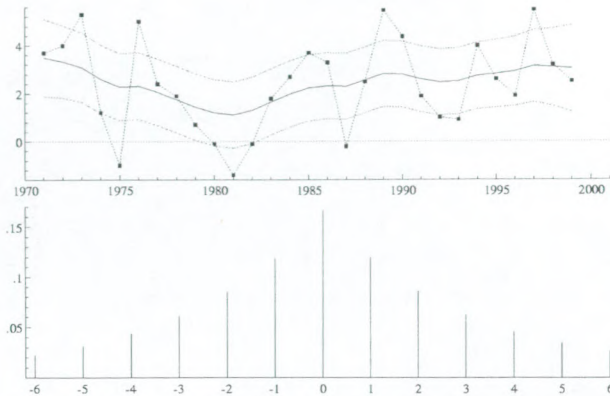


Figure 7: Geschatte curve met 95% betrouwbaarheidsinterval en gewichten ($q = 0.11$)

statistische analyses getoetst kunnen worden op hun effectiviteit. Het is jammer dat niet-parametrische statistici en econometristen het model niet als uitgangspunt nemen. Zij leggen gewichtenpatronen op (deze worden in die literatuur *Kernels* genoemd) zonder dat deze een connectie hebben met een model zodat geen statistische eigenschappen van de geschatte curve afgeleid kunnen worden. Via bepaalde ad-hoc technieken kan men weliswaar onzekerheidsmarges berekenen voor de curve maar dit is kinderspel. Anderen berekenen een curve die gebaseerd is op een zogenaamde tweede of hogere graads polynoom en op het minimaliseren van een som van kwadratische afstanden tussen de observaties en de curve plus een straf-functie voor "gladheid"¹³. Dit is wederom een methode van gladstrijken zonder enige referentie naar een statistisch model. Dat juist statistici en econometristen dit spel spelen is ten minste verrassend te noemen¹⁴. Het is mijn overtuiging dat het model als leidraad zou moeten dienen voor het verantwoord analyseren van waarnemingen.

Ten slotte wil ik opmerken dat het gewichtenpatroon geïmpliceerd door een model

¹³Deze interpolatie methode wordt veelal aangeduid als *spline smoothing* wat kan worden vertaald als gladstrijken met een latje.

¹⁴Een toegankelijk studieboek over niet-parametrische technieken in de econometrie is W. Härdle (1994), "Applied Nonparametric Regression", Cambridge University Press, en die voor in de statistiek is Green en Silverman (1994), "Nonparametric Regression and Generalized Linear Models", London: Chapman & Hall.

een goed hulpmiddel kan zijn om na te gaan of een model effectief is in het gladstrijken van een tijdreeks. In recent werk met Professor Andrew Harvey hebben we aangetoond dat het lokaal nivo model met gecorreleerde storingstermen ε_t en η_t kan leiden tot een asymmetrisch gewichtenpatroon voor het gladstrijken hetgeen onder normale omstandigheden niet wenselijk is ¹⁵.

De nivo tot storing ratio

De connecties tussen de kortingsfaktor λ van de exponentieel gewogen voorspelmethode, de nivo tot storing ratio q van het lokaal nivo model en de coëfficiënt θ van het ARIMA(0,1,1) model heb ik eerder aangegeven. Ik concentreer mij hier op de keuze van een waarde voor de ratio q . De voorspellingen hangen mede af van deze keuze. De meest aannemelijke schatter voor q in ons voorbeeld is eerder uitgerekend. Nu kan het voorkomen dat de zogenaamde aannemelijkheidsfunctie weinig varieert met q , met andere woorden, de aannemelijkheidsfunctie is plat en niet informatief over q . Econometristen concluderen dan veelal dat het model niet goed gedefinieerd is en wijzen een dergelijk model af. Het komt voor dat men met meer vertrouwen zich concentreert op modellen en schattingsmethoden waarin dergelijke problemen niet lijken te bestaan. Ten eerste zou ik willen bestrijden dat modellen dienen te worden afgewezen als deze niet informatief zijn over bepaalde parameters. Zeker, als het coëfficiënten betreft die verbanden moeten blootleggen tussen economische variabelen dan kan men een dergelijk aspect van het model afwijzen zoals dit met behulp van t-toetsen gebeurt in regressie-analyse. Echter het karakter van een coëfficiënt als q is anders dan die van een regressie coëfficiënt. Ratios zoals q worden wel eens denigrerend aangeduid als storingsparameters (in het Engels, *nuisance parameters*). Tot op zekere hoogte zijn het niet de parameters die interessant zijn voor het toetsen van economische theorieën maar ik heb hierboven aangegeven dat de keuze voor q statistisch belangrijk is om te bepalen hoeveel informatie rondom het tijdstip t relevant is om tot een verantwoorde statistische analyse te komen. Wanneer de aannemelijkheidsfunctie weinig informatie geeft over q dan kan men slechts concluderen dat de keuze van q niet zo belangrijk is omdat de informatie in de tijdreeks homogeen over tijd is. Het verwerpen van het model om die reden is buiten proportie.

¹⁵Zie A.C. Harvey en S.J. Koopman (2000) Discussion paper, Vrije Universiteit.

Zoals gezegd, de voorspellingen hangen af van de keuze van een waarde voor q . In de klassieke stroming van de econometrie wordt de onbekende q vervangen door een schatter van q en worden de voorspellingen berekend. De onzekerheidsmarge van de voorspellingen wordt bepaald aan de hand van de kansverdelingsfuncties die zijn opgelegd aan de storingen ε_t en η_t . Met de onzekerheid van de schatter q wordt geen rekening gehouden. Dit is een twee-slachtige houding. We houden rekening met de onzekerheid die geïntroduceerd wordt in het model en welke afhangt van onbekende coëfficiënten maar met de onzekerheid van de schatters van de coëfficiënten wordt geen rekening gehouden. In de Bayesiaanse stroming van de econometrie wordt wel rekening gehouden met deze onzekerheid en dit is een overtuigend aspect van de analyse, zeker wanneer we te maken hebben met slechts weinig waarnemingen. Methoden gebaseerd op simulatie-technieken komen de laatste tien jaar sterk in opmars en laten zich goed lenen voor een Bayesiaanse analyse van het state space model. Het is zeker de intentie om het programma STAMPTM, waarover later meer, aan te passen om een volledige Bayesiaanse analyse uit te kunnen te voeren.

Extensies

De exponentieel gewogen voorspeltechniek is uitgebreid door Holt en Winters voor reeksen met groei en seizoen effecten ¹⁶. De resultaten van Muth zijn door Theil en Wage op een originele manier verder gegeneraliseerd ¹⁷; zij hebben getoond dat de Holt-Winters methode voorspellingen genereert die gelijk zijn aan de minimale verwachting van de kwadratische voorspelfout geïmpliceerd door het lokaal lineaire trend model

$$\begin{aligned}y_t &= \mu_t + \varepsilon_t, \\ \mu_{t+1} &= \mu_t + \beta_t + \eta_t, \\ \beta_{t+1} &= \beta_t + \zeta_t.\end{aligned}$$

Dit model is een extensie van het lokaal nivo model waarin de groeiterm β_t is geïntroduceerd in het tijdsvariërende lokaal nivo model. Bij het tweemaal achtereenvolgens berekenen van de eerste verschillen van de tijdreeks, verkrijgen we het model

$$\Delta^2 y_t = \zeta_{t-1} + \eta_{t-1} - \eta_{t-2} + \varepsilon_t - 2\varepsilon_{t-1} + \varepsilon_{t-2}.$$

¹⁶Zie C.C. Holt (1957), Research Memorandum 52, Carnegie Institute of Technology, Pittsburgh, Pennsylvania, USA, en P.R. Winters (1960), Management Science 6, p.324-342.

¹⁷Zie H. Theil en S. Wage (1964), Management Science 10, p.198-206.

De auto-correlaties voor vertraging 1 en 2 perioden zijn ongelijk aan nul terwijl alle andere auto-correlaties gelijk zijn aan nul. Een reeks gegenereerd door dit model heeft dezelfde eigenschappen als een ARIMA(0,2,2) model.

We zien dat twee verschillende modellen, welke conceptueel zeer verschillend van elkaar zijn, dezelfde voorspellingen kunnen genereren. De exponentieel gewogen voorspelmethode kan gebaseerd zijn op een state space model of op een Box-Jenkins model. Het toepassen van het principe van de minimum verwachting van de kwadratische voorspelfout voor beide modellen resulteert in een variant van de exponentieel gewogen voorspelmethode. Dit leidt ons naar de conclusie dat, wanneer de tijdreeks een onderliggende structuur heeft die eenvoudig genoeg is, de twee methodologiën in essentie gelijk zijn. Het is wanneer we meer complexe situaties bestuderen dat de twee methodologiën van elkaar divergeren. In het vervolg van deze rede zal ik dit nader onderstrepen. Ik zal beginnen met een formeel betoog waarin ik de twee methodologiën met elkaar vergelijk.

State space model versus Box-Jenkins model

De eerste ontwikkelingen van de state space methodologie hebben plaats gevonden in de technische wetenschappen in plaats van de statistiek of econometrie en vinden hun oorsprong in de invloedrijke paper van Professor Kalman ¹⁸. In deze paper doet Kalman twee belangrijke dingen. Hij laat zien dat een zeer brede klasse van modellen en problemen kunnen worden beschreven door een simpel lineair model welke in essentie gelijk is aan het state space model zoals eerder geformuleerd is. Hij liet verder zien dat de berekeningen voor het analyseren van tijdreeksen op basis van het state space model kunnen worden uitgevoerd op een recursieve manier die zeer geschikt is voor implementatie op een computer. Heel veel werk is er vervolgens verricht in de ontwikkeling van deze ideeën in de technische literatuur.

De contributies van statistici en econometristen in de jaren zestig, zeventig en tachtig op het gebied van de state space methodologie waren summier en sporadisch. In deze rede zal ik geen volledig historisch overzicht geven van de verschillende contributies. Maar de namen Jones, Akaike, Duncan, Harrison, Stevens, Kitagawa, Ansley, Kohn en de Jong zijn belangrijk in deze. In het bijzonder wil ik mijn leermeester

¹⁸Zie R.E. Kalman (1960), Journal of Basic Engineering, Transactions ASME, Series D 82, p.35-45.

Professor Andrew Harvey noemen, niet alleen voor zijn wetenschappelijke contributies maar ook voor zijn boek uit 1989 en zijn initiatieven voor het computerprogramma STAMPTM ¹⁹. Zijn boek beschrijft de ontwikkelingen in de state space tijdreeksanalyse tot de tweede helft van de jaren tachtig. Hij plaatst ook de ontwikkelingen van state space binnen de econometrie in een historisch perspectief.

Het grote voordeel van state space is dat deze gebaseerd is op een structurele analyse van het tijdreeksprobleem. De verschillende componenten waaruit de tijdreeks bestaat, zoals trend, seizoen, cyclische bewegingen en kalender effecten, tezamen met exogene verklarende variabelen en interventies, kunnen apart gemodelleerd worden voordat zij aaneen worden geschakeld in een state space model om daarna simultaan te kunnen worden geanalyseerd. Het is vrij aan de gebruiker van het model om specifieke componenten van een tijdreeks in een model te introduceren en vervolgens te identificeren met behulp van beschikbare observaties. Ter vergelijking, de Box-Jenkins methodologie is een zwarte doos. Het gekozen model hangt af van vage keuzen voor het aantal eenheidswortels in het autoregressieve gedeelte van het ARIMA model en voor de maximale relevante vertraging in het model. De Box-Jenkins methodologie levert vervolgens schattingen van coëfficiënten op die weinig informatie verschaft aan zowel de beginnende als de gevorderde gebruiker.

Een ander voordeel van het state space model is dat deze zeer flexibel is. De recursieve structuur van zowel het model als de verwante recursieve methoden en technieken, maakt het eenvoudig om veranderingen in de structuur van het model te introduceren over tijd. Bijvoorbeeld het combineren van maandelijks observaties met kwartaal observaties uit een eerdere periode is relatief eenvoudig te modelleren in een state space model. Dit in tegenstelling tot het Box-Jenkins model dat homogeen over tijd dient te zijn omdat de aanverwante schattings-techniek afhangt van de assumptie dat de tijdreeks, in eerste of tweede verschillen, stationair dient te zijn. Dit betekent dat de dynamische eigenschappen van de gehele tijdreeks, voor elke periode, gelijk dient te zijn. Dit is in de praktijk een zeer stringente assumptie.

State space modellen zijn zeer algemeen: een lange reeks van lineaire modellen, inclusief de ARIMA modellen van Box en Jenkins, kunnen worden gerepresenteerd als een state space model. Multivariate modellen in state space zijn relatief eenvoudige

¹⁹Zie S.J. Koopman, A.C. Harvey, J.A. Doornik en N. Shephard (2000), London: Timberlake Consultants.

extensies van univariate modellen en de schattingstechnieken hoeven niet wezenlijk te worden aangepast hetgeen niet het geval is voor de Box-Jenkins methodologie.

Met state space technieken is het makkelijk rekening te houden met niet waargenomen observaties in de reeks hetgeen in de praktijk kan voorkomen door storingen in bedrijfsprocessen of door onbeholpenheid. Verder, regressie coëfficiënten kunnen stochastisch variëren over tijd wanneer dat in bepaalde omstandigheden wenselijk wordt geacht. Geen extra theorie is nodig voor het voorspellen van tijdreeksen. Alles wat nodig is, is het projecteren van het Kalman filter op toekomstige niet-waargenomen observaties. De technieken die we eerder hebben gebruikt voor het analyseren van de beschikbare observaties, zijn ook in staat voorspellingen te produceren tezamen met betrouwbaarheidsintervallen.

U kunt zich nu afvragen, als dit allemaal de voordelen zijn van state space modellen, wat zijn dan de nadelen in vergelijking met de methodologie van Box en Jenkins? Naar mijn stellige overtuiging zijn er slechts twee nadelen te noemen:

- (i) het gemis binnen de econometrische en statistische gemeenschap van informatie en kennis over state space;
- (ii) de weinig beschikbare computerprogrammatuur voor analyses gebaseerd op state space modellen.

De Box-Jenkins methodologie is een wezenlijk onderdeel van veel cursussen in de statistiek en econometrie op Amerikaanse en Europese universiteiten. Daardoor worden de Box-Jenkins methoden veelal opgenomen in de studieboeken voor tijdreeksen en econometrie. Verder hebben de standaard computerprogramma's voor econometrische en statische analyses, zoals SASTM, TSPTM en EvIEWSTM, verschillende opties beschikbaar voor een Box-Jenkins analyse. Daarentegen wordt aan relatief weinig universiteiten de state space methodologie opgenomen in het basisprogramma voor de studie economie of econometrie. In tegenstelling tot de technische wetenschappen, zijn er slechts een aantal studieboeken met hoofdstukken over state space waarvan die van Professor Harvey wel het meest volledige is. Hierbij dient te worden aangetekend dat in een aantal recente econometrische studieboeken de state space methoden wel worden behandeld. Verder zijn er ook relatief weinig computerprogramma's die de methoden en technieken zoals het Kalman filter in zijn algemeenheid ondersteunen. Ik kom later terug op bestaande computerprogramma's die wel state space methoden gebruiken zoals STAMPTM en OxTM/SsfPackTM.

Ik ga nu in op een aantal nadelen van de Box-Jenkins methoden. Het elimineren van trend en seizoen effecten door de data te transformeren, via het nemen van verschillen, hoeft op zichzelf geen nadeel te zijn als het enigste doel van de analyse het berekenen van voorspellingen is. Maar in veel gevallen, en zeker in de econometrie, hebben de verschillende componenten in een state space model een intrinsieke betekenis die van belang kan zijn. Natuurlijk, in een Box-Jenkins analyse kunnen de componenten geres- taureerd worden door een soort van retour transformatie die gebaseerd is op het max- imaliseren van de verwachting van geconstrueerde kwadratische residuen. Maar dit kan worden afgedaan als een kunstmatige en weinig aantrekkelijke handreiking.

De eis dat na het nemen van verschillen de data stationair zou moeten zijn is een zwakke plek in de theorie. In de economie, maar ook in andere sociale wetenschappen, zijn geobserveerde reeksen zoals het nationaal inkomen en de consumptie nooit station- air, hoeveel keer er ook verschillen worden genomen. Dus het is aan de onderzoeker om te bepalen hoe nabij stationairiteit benaderd is. Maar hoe dichtbij is dichtbij genoeg ? Dit is een zeer moeilijke vraag om te beantwoorden en het heeft geleid tot een reeks van zogenaamde eenheidswortel toetsen die in een laboratorium (een simulatie-studie) soms goed voor de dag komen maar die zich empirisch nog een weinig hebben bewezen.

In de praktijk wordt veelal geconcludeerd dat specifieke Box-Jenkins ARIMA mod- ellen zoals het zogenaamde "airline" model, goed zijn in het voorspellen van tijdreeksen maar ik heb eerder gezegd dat deze modellen ongeveer gelijk zijn aan plausibele state space modellen. Het is relatief eenvoudig te laten zien dat het lokaal nivo model gelijk is aan een $ARIMA(0,1,1)$ model met een geresliceerde waarde voor de coëfficiënt θ . Het- zelfde geldt voor modellen met trend, seizoen en andere componenten. De restricties op de coëfficiënten kunnen ervoor zorgen dat niet het optimale model wordt gebruikt om de voorspellingen te genereren. Ik heb echter een voorkeur voor voorspellingen die berek- end worden met behulp van modellen die de eigenschappen van de tijdreeks zo effectief mogelijk beschrijven. Voor tijdreeksprocessen, die niet alleen bestaan uit trend, seizoen en storingstermen, is het identificeren van een Box-Jenkins model veel moeilijker. Het enigste identificatie hulpmiddel is de steekproef auto-correlatie functie welke bekend is om zijn onnauwkeurigheid dat veroorzaakt wordt door de hoge steekproef variabiliteit. Het kan daardoor gemakkelijk voorkomen dat verschillende ARIMA modellen worden geïdentificeerd die dezelfde tijdreeks op even nauwkeurige wijze kunnen reproduceren. Op basis van de redenen hier eerder aangegeven is de state space methodologie in staat

de componenten van een tijdreeks effectiever te identificeren dan binnen de Box-Jenkins methodologie.

State space in de econometrie

Seizoencorrectie

Naast het voorspellen, is het corrigeren van tijdreeksen voor seizoenfluctuaties veruit één van de meest gebruikte toepassingen van tijdreeksanalyse. Het doel van seizoencorrectie is het corrigeren van tijdreeksen voor seizoenseffecten. In de meeste statistische bureaus en elders wordt dit gedaan door het toepassen van de zogenaamde X11 methode van het Amerikaanse Census bureau. Vele tijdreeks-specialisten binnen de universitaire wereld zijn hier verrast over, zeker wanneer men dergelijke technieken nader bestudeert. Men zou eerder verwachten dat een methode gebruikt zou worden welke gebaseerd is op een model. De basistechnieken achter de X11 methode zijn zo'n dertig jaar geleden ontwikkeld en zijn niet essentieel veranderd in die tijd. Wat dit betreft dienen we deze Amerikaanse statistici te feliciteren met dit grote succes dat stand houdt zelfs na de revolutionaire ontwikkelingen in de computertechnologie hetgeen het mogelijk maakt zelfs ingewikkelde modellen snel te schatten.

Een deel van de verklaring van het succes is dat trend en seizoen niet geobserveerd kunnen worden en dus kunnen we nooit echt weten of het seizoen correct geschat is en of een andere methode het beter kan. Hoe glad moet een trend of een seizoen er uit zien? Er bestaat geen objectief criterium voor het antwoord op deze vraag²⁰. De Census X11 ontwerpers hebben op de werkvloer, waarin men door gemaakte fouten wijzer kon worden, een methode ontwikkeld welke seizoenseffecten bepaalt die in de meeste gevallen voldoet voor de meeste gebruikers.

Desalniettemin, de X11 methode, de latere X12 methode en varianten ervan zijn verre van perfect, zeker vanuit een statistisch oogpunt. Bijvoorbeeld, deze methoden maken niet optimaal gebruik van de observaties aan het einde van de tijdreeks²¹. Met

²⁰Professoren den Butter (VU) en Fase (UvA, DNB) hebben in 1991 het boek "Seasonal Adjustment as a practical problem" geschreven waarin een aantal criteria worden gegeven waaraan seizoenscorrectie methoden zouden moeten voldoen.

²¹De ontwerpers van X11 zijn hiervan op de hoogte en men verweert zich met de wens om zogenaamde revisies te minimaliseren.

de verdere ontwikkelingen van state space methoden, de groeiende beschikbaarheid van informatie met daaraan gekoppeld een groeiende vraag naar voorspellingen en seizoensgecorrigeerde reeksen, geloof ik dat er meer aandacht zal komen voor methoden welke gegrond zijn op statistische modellen die optimaal gebruik maken van de beschikbare informatie. Verklarende variabelen, met name dagkarakteristieken, kalender- effecten en weergegevens, zijn relatief eenvoudig te incorporeren in state space methoden.

Het is mijn intentie om in dialoog te blijven met de makers en gebruikers van X11 en hen te overtuigen van een state space benadering van het probleem. Er is op dit terrein veel werk te verrichten en ik ben van plan om te blijven deelnemen aan verdere ontwikkelingen met name voor het corrigeren van observaties die wekelijks, dagelijks of per uur worden waargenomen.

Econometrische modellen

Het introduceren van verklarende variabelen in het state space tijdreeksmodel is eenvoudig. Het lokaal nivo model kan bijvoorbeeld als volgt worden uitgebreid

$$y_t = \mu_t + x_t' \beta + \varepsilon_t, \quad \mu_{t+1} = \mu_t + \eta_t,$$

waarbij x_t een vector is van verklarende variabelen en β een vector van onbekende coëfficiënten. De constante van het klassieke lineaire regressie model is hier vervangen door een tijdsvariërende constante. In veel situaties kan het lokaal nivo in dit model geïnterpreteerd worden als sociale, maatschappelijke of institutionele invloeden welke niet eenvoudig te kwantificeren zijn. De onbekende coëfficiënten in de vector β kunnen, gelijk met het lokale nivo component μ_t , via het Kalman filter geschat worden. Het Kalman filter is vanuit dit gezichtspunt gelijk aan de kleinste kwadraten methode maar dan voor regressie modellen met tijdsvariërende coëfficiënten zoals μ_t . De coëfficiënten vector β kunnen we ook laten variëren over tijd. In speciale gevallen kan dit een nuttige extensie zijn maar in de meeste gevallen is dit ongewenst. Economische theorieën zijn geschoeid op stabiele relaties tussen economische grootheden en wanneer deze flexibel worden gemaakt kan dit wellicht tot een beter schattingsresultaat leiden maar het werkt het identificeren van schijn-relaties in de hand en dit verscherpt ons economisch inzicht geenszins. Ik wil hier graag onderstrepen dat state space methoden prima hand in hand kunnen gaan met andere modellen, methoden en technieken in de tijdreeks-econometrie.

Dynamische interacties tussen economische variabelen kunnen effectief worden gemodelleerd door een multivariaat model. Binnen de state space methodologie kan men op eenvoudige wijze eerst modellen bepalen voor de individuele reeksen en vervolgens de modellen bij elkaar nemen om de interacties tussen de reeksen te bestuderen. In de volgende illustratie nemen we twee tijdreeksen y_t en z_t die de karakteristieken hebben van een lokaal nivo model:

$$\begin{aligned}y_t &= \mu_{y,t} + \varepsilon_{y,t}, & \mu_{y,t+1} &= \mu_{y,t} + \eta_{y,t}, \\z_t &= \mu_{z,t} + \varepsilon_{z,t}, & \mu_{z,t+1} &= \mu_{z,t} + \eta_{z,t},\end{aligned}$$

waarbij $\mu_{y,t}$ and $\mu_{z,t}$ het nivo, $\varepsilon_{y,t}$ en $\varepsilon_{z,t}$ de storing en $\eta_{y,t}$ en $\eta_{z,t}$ de verandering van het nivo representeren voor de tijdreeksen y_t en z_t . We kunnen de twee tijdreeksen y_t en z_t simultaan modelleren door correlatie coëfficiënten ρ_ε , voor de storingen $\varepsilon_{y,t}$ en $\varepsilon_{z,t}$, en ρ_η , voor de nivo-veranderingen $\eta_{y,t}$ en $\eta_{z,t}$, te introduceren welke de waarden tussen -1 en 1 kunnen aannemen. Wanneer de correlatie ρ_η gelijk is aan 1 of -1 dan kunnen we het nivo van z_t schrijven als een lineaire combinatie van het nivo van y_t :

$$\mu_{z,t} = a + b\mu_{y,t},$$

waarbij a en b onbekende coëfficiënten zijn. Hieruit volgt dat de tijdreeks

$$by_t - z_t,$$

stationair is hetgeen betekent dat y_t en z_t zogenaamd gecointegreerd zijn. Deze benadering leidt tot een alternatieve methodologie voor het bestuderen van stabiele dynamieke relaties tussen economische variabelen hetgeen de afgelopen tien jaar veel aandacht in de literatuur heeft gehad. In de empirische tijdreeks econometrie van vandaag de dag maakt men veelvuldig gebruik van multivariate versies van het autoregressieve model. Het debat over deze benadering of die op basis van het state space model heeft een aantal gelijke kenmerken van de discussie die we hebben gevoerd over de relaties tussen het Box-Jenkins en het state space model. Ik wil hierbij onderstrepen dat de state space benadering niet afhangt van zogenaamde eenheidswortel toetsen welke een hoge steekproef variabiliteit kennen.

Het bovenstaande laat zien dat state space tijdreeksmodellen gemakkelijk de stap naar econometrische modellen maken waarin verbanden kunnen worden gelegd tussen economische variabelen. In het ene geval wordt op eenvoudige wijze exogene variabelen

geïntroduceerd in het state space tijdreeksmodel. In het andere geval worden de methoden die gebruikt worden voor univariate tijdreeksen toegepast op multivariate tijdreeksen. Wanneer we de state space methodologie gebruiken voor multivariate tijdreeksen, waarvan ik zojuist een eenvoudig voorbeeld heb gegeven, dan ontstaat er een zeer algemene klasse van econometrische modellen. In mijn optiek zijn het de redenen die ik eerder heb genoemd in de discussie over state space versus Box-Jenkins analyse, die verklaren dat state space methoden nog weinig gebruikt worden in de literatuur. Ik voel het als een van mijn opdrachten om verder empirisch georiënteerd werk uit te voeren met behulp van multivariate state space tijdreeksmodellen.

Modellen gebaseerd op andere distributies

Tot dusver in deze lezing heb ik mij gericht op standaard lineaire modellen waarbij storings-termen uit een normale kansverdeling zijn gegenereerd. Deze modellen zijn in mijn ogen al complex genoeg voor effectief gebruik in de empirische econometrie, dat er daardoor weinig redenen bestaan om modellen verder te compliceren zeker als de statistische eigenschappen van dergelijke modellen onbekend zijn. Desalniettemin, tijdreeksen kunnen intrinsiek niet-lineair zijn of kunnen gegenereerd zijn door discrete kansverdelingen zodat het moeilijk is vast te houden aan het zogenaamde lineaire normale model. Ik zal twee voorbeelden geven.

- Het modelleren van onzekerheden op financiële en andere intensieve handelsmarkten is een actief onderzoeksterrein in de econometrie. Een mogelijk model voor financiële opbrengsten is gegeven door

$$\begin{aligned} y_t &= \mu + \varepsilon_t, & \text{Var}(\varepsilon_t) &= \exp(h_t), \\ h_{t+1} &= \phi h_t + \eta_t, & \text{Var}(\varepsilon_t) &= \sigma^2, \quad |\phi| < 1, \end{aligned}$$

waarbij y_t een indicatie is van een gerealiseerd rendement op een financieel fonds in een bepaalde tijdsperiode hetgeen moeilijk te voorspellen is en daarom slechts wordt gemodelleerd als een constante μ plus een storing ε_t met verwachting nul. De variantie laten we variëren over tijd door de log-transformatie te modelleren als een autoregressief proces. Het volgt uit de formulering van deze vereenvoudigde versie van het zogenaamde *stochastisch volatiliteits model* dat de waarneming y_t niet lineair is in het tijdreeksproces h_t zoals y_t wel lineair is in μ_t van het lokaal nivo model. We hebben daardoor met een niet-lineair model van doen.

- Het modelleren van variabelen die discreet of kwalitatief worden waargenomen is een ander actief onderzoeksterrein en welke belangrijk is voor de empirische econometrie. Een voorbeeld is een tijdreeks van het aantal voltooide transacties op een dag hetgeen kan variëren tussen nul en een bepaald maximum, zeg twintig. Een mogelijk model voor dit soort discrete gegevens is

$$y_t \sim \text{Poisson}[\exp(\mu_t)], \quad \mu_{t+1} = \mu_t + \eta_t.$$

Het is moeilijk om hier y_t te modelleren als een lineair normaal model.

Er zijn de laatste jaren verschillende methoden en technieken ontwikkeld voor het schatten van dergelijke niet-standaard tijdreeksmodellen van zowel een klassiek statistische invalshoek als wel een Bayesiaanse invalshoek. De methodologie is gestoeld op simulatie-technieken die het mogelijk maken om de meest aannemelijkheidsfunctie en schatters voor de state vector te evalueren tezamen met conditionele (of posteriore) kansverdelingsfuncties²². Het werk dat ik op dit gebied heb gedaan in samenwerking met Professor James Durbin wordt uitvoerig beschreven in deel II van ons eerder genoemde boek. Het vergt te veel tijd en technische achtergrond om in deze rede verder in te gaan op deze technieken.

Computerprogrammatuur

Eén van mijn andere activiteiten in dit gebied van de wetenschap is het ontwikkelen van technische computerprogrammatuur voor state space methoden, zoals het Kalman filter, met als uiteindelijk doel de geavanceerde technieken op een gebruikersvriendelijke wijze aan te bieden. Zoals ik eerder heb opgemerkt, een van de redenen dat state space methoden nog weinig gebruikt worden in de praktijk is gelegen in het feit dat deze niet beschikbaar zijn als een optie in de standaard software-pakketten voor tijdreeksanalyse²³. Het is de algemene perceptie dat het niet eenvoudig is om het Kalman filter voor een bepaald model correct te implementeren op een computer. Tot op zekere hoogte

²²Zie N. Shephard en M.K. Pitt (1997), *Biometrika* 84, p.653-667, J. Durbin en S.J. Koopman (1997), *Biometrika* 84, p.669-684, en J. Durbin en S.J. Koopman (2000), *Journal of the Royal Statistical Society, Series B* 62, p.3-56

²³Ten tijde van het reviseren van deze rede voor de KM publicatie, hebben de makers van het econometrische software pakket EvIEWSTM aangekondigd spoedig met een module te komen voor state space modellen gelijkend op SsfPackTM.

is dit waar en dient rekening gehouden te worden met een investering in programmeertijd. Het computerprogramma STAMPTM, waarbij ik sinds 1993 betrokken ben samen met Andrew Harvey, Jurgen Doornik en Neil Shephard, is een voorbeeld van een gebruikersvriendelijk programma dat een volledige state space analyse kan uitvoeren voor een bepaalde klasse van univariate en multivariate tijdreeksmodellen ²⁴. De gebruiker verliest geen tijd aan het programmeren en er zijn genoeg opties beschikbaar om een tijdreeks op een effectieve wijze te analyseren.

Voor het algemene state space model is het moeilijker een algemeen gebruikersvriendelijk computerprogramma te ontwikkelen omdat de technieken op een flexibele wijze moeten worden aangeboden. Daarnaast willen collega wetenschappers meer flexibiliteit wanneer de technieken gebruikt worden voor het oplossen van onderzoeksvraagstukken. Het is daardoor beter om de methoden en technieken aan te bieden als een onderdeel in bestaande econometrische en statistische programmeertalen zoals OxTM, SplusTM en MatlabTM. De software onderdelen voor het algemene state space model zijn verzameld in de programma-bibliotheek SsfPackTM en samen ontwikkeld met Neil Shephard en Jurgen Doornik ²⁵. In dit geval dient men enige tijd aan het programmeren te besteden maar min of meer alleen voor het formuleren van het model en de presentatie van de resultaten van de analyses.

Het is mijn intentie om verdere ontwikkelingen in de state space methodologie te volgen en zelf te initiëren. Met dezelfde energie zal ik de computerprogramma's STAMP en SsfPack verder ontwikkelen zodat de laatste ontwikkelingen beschikbaar komen op een gebruikersvriendelijke wijze voor een ieder die deze methoden wil gebruiken.

²⁴STAMP is een afkorting van *Structural Time series Analyser, Modeller and Predictor*. Het wordt uitgegeven door Timberlake Consultants te Londen en maakt deel uit van de OxMetricsTM groep van econometrische software. Het is geschreven voor zogenaamde *structurele tijdreeksmodellen* waarvan het lokaal nivo model het meest eenvoudige voorbeeld is. Voor meer informatie: <http://www.ssfpack.com/stamp/>.

²⁵Zie Koopman, Shephard en Doornik (1999), *Econometrics Journal* 2, p.113-166. Voor meer informatie: <http://www.ssfpack.com>.

Voorspelling van de groei van het nationaal inkomen voor 2001

Het geeft pas, bij het breken van een lans voor de state space methodologie in tijdreeks-analyse, om de daad bij het woord te voegen: sta op, en voorspel een tijdreeks. Eerder in deze rede heb ik de groei van het nationaal inkomen gebruikt om een aantal aspecten van de verschillende methoden te illustreren. De observaties waren beschikbaar van 1970 tot en met 1999. Het Centraal Planbureau (CPB) te Den Haag heeft de groei voor dit jaar, 2000, voorspelt op 3.25%. Op basis van het lokaal nivo model en de jaarlijkse observaties voor 1970-1999, komt mijn voorspelling uit op 3.05% met een 50% betrouwbaarheidsmarge van $\pm 1.33\%$. Wat dit betreft komt de CPB voorspelling goed overeen met de voorspelling van de relatief eenvoudige exponentionele gewichten-methode. Ik ga nog een stap verder. Ervan uitgaande dat het CPB gelijk zal krijgen met een groei van 3.25% dit jaar, dan voorspel ik op basis van het lokaal nivo model dat de groei voor volgend jaar, 2001, terecht zal komen op 3.10% met wederom een 50% betrouwbaarheidsmarge van zeg $\pm 1.33\%$. Men mag mij eraan houden ! ²⁶.

²⁶Een dergelijke belofte is makkelijk gedaan gegeven het feit dat recentelijk vele economen, en niet de minste onder hen, moesten terug komen op eerder gedane uitspraken over de invloed van de Azië crisis op de groei van de Westerse economiën.

