

Makroekonomia Zaawansowana

Ćwiczenia 9 – Estymacja modeli DSGE za pomocą metod bayesowskich

Andrzej Torój



Instytut Ekonometrii |
Zakład Ekonometrii Stosowanej

Plan ćwiczeń

- 1 Estymacja bayesowska
- 2 Rozkłady a priori
- 3 Konstrukcja funkcji wiarygodności
- 4 Rozkład a posteriori
- 5 W praktyce

Plan prezentacji

- 1 Estymacja bayesowska
- 2 Rozkłady a priori
- 3 Konstrukcja funkcji wiarygodności
- 4 Rozkład a posteriori
- 5 W praktyce

Ekonometria bayesowska – naczelna idea

- w klasycznej ekonometrii:
 - cała wiedza o wartości parametru pochodzi z próby statystycznej
 - szacujemy parametr punktowo, ew. z przedziałem ufności
- w bayesowskiej ekonometrii:
 - parametr traktujemy nie jako jedną “prawdziwą” wartość, ale jako zmienną losową (tzn. może przyjmować wiele różnych wartości z określonym prawdopodobieństwem – otrzymujemy ROZKŁAD)
 - wiedza o wartości parametru pochodzi częściowo z próby, ale częściowo jest narzucona z góry (tzw. rozkład *a priori*)

Dlaczego metody bayesowskie w DSGE?

Rozwiązany zlinearyzowany model DSGE jest przecież postaci $y_t = My_{t-1} + Nf_t$ i można go szacować metodą największej wiarygodności tak jak model SVAR...

- silne restrykcje teoretyczne dla ekonomicznie interpretowalnych wartości parametrów
- ucieczka parametrów w bezsensowne ekonomicznie zakresy grozi złamaniem warunków Blancharda-Kahna
- w badaniach empirycznych często dość krótkie próby – potrzeba informacji *a priori*
- model DSGE jest na ogół zbyt silnie wyidealizowany (zawiera zbyt mało sprzężeń w porównaniu z rzeczywistością), aby bez “wsparcia” *a priori* być skonfrontowanym z danymi (np. w ramach metody FIML)
- wiedza spoza próby statystycznej jest wprowadzana w sposób systematyczny, a nie nieformalny
- czasami mamy w modelu więcej zmiennych obserwowalnych niż wstrząsów – przy metodach klasycznych to problem

Dlaczego metody bayesowskie w DSGE?

Rozwiązany zlinearyzowany model DSGE jest przecież postaci $y_t = My_{t-1} + Nf_t$ i można go szacować metodą największej wiarygodności tak jak model SVAR...

- silne restrykcje teoretyczne dla ekonomicznie interpretowalnych wartości parametrów
- ucieczka parametrów w bezsensowne ekonomicznie zakresy grozi złamaniem warunków Blancharda-Kahna
- w badaniach empirycznych często dość krótkie próby – potrzeba informacji *a priori*
- model DSGE jest na ogół zbyt silnie wyidealizowany (zawiera zbyt mało sprzężeń w porównaniu z rzeczywistością), aby bez “wsparcia” *a priori* być skonfrontowanym z danymi (np. w ramach metody FIML)
- wiedza spoza próby statystycznej jest wprowadzana w sposób systematyczny, a nie nieformalny
- czasami mamy w modelu więcej zmiennych obserwowalnych niż wstrząsów – przy metodach klasycznych to problem

Dlaczego metody bayesowskie w DSGE?

Rozwiązany zlinearyzowany model DSGE jest przecież postaci $y_t = My_{t-1} + Nf_t$ i można go szacować metodą największej wiarygodności tak jak model SVAR...

- silne restrykcje teoretyczne dla ekonomicznie interpretowalnych wartości parametrów
- ucieczka parametrów w bezsensowne ekonomicznie zakresy grozi złamaniem warunków Blancharda-Kahna
- w badaniach empirycznych często dość krótkie próby – potrzeba informacji *a priori*
- model DSGE jest na ogół zbyt silnie wyidealizowany (zawiera zbyt mało sprzężeń w porównaniu z rzeczywistością), aby bez “wsparcia” *a priori* być skonfrontowanym z danymi (np. w ramach metody FIML)
- wiedza spoza próby statystycznej jest wprowadzana w sposób systematyczny, a nie nieformalny
- czasami mamy w modelu więcej zmiennych obserwowalnych niż wstrząsów – przy metodach klasycznych to problem

Dlaczego metody bayesowskie w DSGE?

Rozwiązany zlinearyzowany model DSGE jest przecież postaci $y_t = My_{t-1} + Nf_t$ i można go szacować metodą największej wiarygodności tak jak model SVAR...

- silne restrykcje teoretyczne dla ekonomicznie interpretowalnych wartości parametrów
- ucieczka parametrów w bezsensowne ekonomicznie zakresy grozi złamaniem warunków Blancharda-Kahna
- w badaniach empirycznych często dość krótkie próby – potrzeba informacji *a priori*
- model DSGE jest na ogół zbyt silnie wyidealizowany (zawiera zbyt mało sprzężeń w porównaniu z rzeczywistością), aby bez “wsparcia” *a priori* być skonfrontowanym z danymi (np. w ramach metody FIML)
- wiedza spoza próby statystycznej jest wprowadzana w sposób systematyczny, a nie nieformalny
- czasami mamy w modelu więcej zmiennych obserwowalnych niż wstrząsów – przy metodach klasycznych to problem

Dlaczego metody bayesowskie w DSGE?

Rozwiązany zlinearyzowany model DSGE jest przecież postaci $y_t = My_{t-1} + Nf_t$ i można go szacować metodą największej wiarygodności tak jak model SVAR...

- silne restrykcje teoretyczne dla ekonomicznie interpretowalnych wartości parametrów
- ucieczka parametrów w bezsensowne ekonomicznie zakresy grozi złamaniem warunków Blancharda-Kahna
- w badaniach empirycznych często dość krótkie próby – potrzeba informacji *a priori*
- model DSGE jest na ogół zbyt silnie wyidealizowany (zawiera zbyt mało sprzężeń w porównaniu z rzeczywistością), aby bez “wsparcia” *a priori* być skonfrontowanym z danymi (np. w ramach metody FIML)
- wiedza spoza próby statystycznej jest wprowadzana w sposób systematyczny, a nie nieformalny
- czasami mamy w modelu więcej zmiennych obserwowalnych niż wstrząsów – przy metodach klasycznych to problem

Dlaczego metody bayesowskie w DSGE?

Rozwiązany zlinearyzowany model DSGE jest przecież postaci $y_t = My_{t-1} + Nf_t$ i można go szacować metodą największej wiarygodności tak jak model SVAR...

- silne restrykcje teoretyczne dla ekonomicznie interpretowalnych wartości parametrów
- ucieczka parametrów w bezsensowne ekonomicznie zakresy grozi złamaniem warunków Blancharda-Kahna
- w badaniach empirycznych często dość krótkie próby – potrzeba informacji *a priori*
- model DSGE jest na ogół zbyt silnie wyidealizowany (zawiera zbyt mało sprzężeń w porównaniu z rzeczywistością), aby bez “wsparcia” *a priori* być skonfrontowanym z danymi (np. w ramach metody FIML)
- wiedza spoza próby statystycznej jest wprowadzana w sposób systematyczny, a nie nieformalny
- czasami mamy w modelu więcej zmiennych obserwowalnych niż wstrząsów – przy metodach klasycznych to problem

Podstawowy wzór ekonometrii bayesowskiej

Łączne prawdopodobieństwo danych i parametrów:

$$P(Y, \theta) = \begin{cases} P(Y|\theta) \cdot P(\theta) \\ P(\theta|Y) \cdot P(Y) \end{cases}$$

Stąd reguła Bayesa:

$$P(\theta|Y) = \frac{P(Y|\theta)}{P(Y)} \cdot P(\theta)$$

Innymi słowy, traktujemy tu wektor parametrów θ jak wielowymiarową zmienną losową. Jej rozkład warunkowy względem danych $P(\theta|Y)$ zależy od rozkładu bezwarunkowego parametrów $P(\theta)$ oraz rozkładu danych warunkowego względem parametrów $P(Y|\theta)$.

Rozkład łączny *a posteriori*

$$P(\theta|Y) = \frac{P(Y|\theta)}{P(Y)} \cdot P(\theta)$$

- $P(\theta)$ to gęstość *a priori*, obrazująca wiedzę o parametrze posiadaną zanim przystąpiliśmy do estymacji. O tym rozkładzie decydujemy sami w taki sposób, by odzwierciedlić nasze założenia, np. parametr ma należeć do przedziału od 0 do 1 (czyli przypisujemy prawdopodobieństwo *a priori* zero innym wartościom rzeczywistym) albo parametr ma wynosić około 0,6 (wówczas ustalamy taką średnią i niską wariancję wokół niej).
- $P(Y|\theta)$ to gęstość próbkowa (identyczna z funkcją wiarygodności wykorzystywaną w metodach ML, FIML itp.).
- $P(Y)$ często pomijamy jako stałą skalującą (gwarantującą że funkcja gęstości całkuje się do 1), zapisując proporcjonalność $P(\theta|Y) \propto P(Y|\theta) \cdot P(\theta)$.

Rozkład łączny *a posteriori*

$$P(\theta|Y) = \frac{P(Y|\theta)}{P(Y)} \cdot P(\theta)$$

- $P(\theta)$ to gęstość *a priori*, obrazująca wiedzę o parametrze posiadaną zanim przystąpiliśmy do estymacji. O tym rozkładzie decydujemy sami w taki sposób, by odzwierciedlić nasze założenia, np. parametr ma należeć do przedziału od 0 do 1 (czyli przypisujemy prawdopodobieństwo *a priori* zero innym wartościom rzeczywistym) albo parametr ma wynosić około 0,6 (wówczas ustalamy taką średnią i niską wariancję wokół niej).
- $P(Y|\theta)$ to gęstość próbkowa (identyczna z funkcją wiarygodności wykorzystywaną w metodach ML, FIML itp.).
- $P(Y)$ często pomijamy jako stałą skalującą (gwarantującą że funkcja gęstości całkuje się do 1), zapisując proporcjonalność $P(\theta|Y) \propto P(Y|\theta) \cdot P(\theta)$.

Rozkład łączny *a posteriori*

$$P(\theta|Y) = \frac{P(Y|\theta)}{P(Y)} \cdot P(\theta)$$

- $P(\theta)$ to gęstość *a priori*, obrazująca wiedzę o parametrze posiadaną zanim przystąpiliśmy do estymacji. O tym rozkładzie decydujemy sami w taki sposób, by odzwierciedlić nasze założenia, np. parametr ma należeć do przedziału od 0 do 1 (czyli przypisujemy prawdopodobieństwo *a priori* zero innym wartościom rzeczywistym) albo parametr ma wynosić około 0,6 (wówczas ustalamy taką średnią i niską wariancję wokół niej).
- $P(Y|\theta)$ to gęstość próbkowa (identyczna z funkcją wiarygodności wykorzystywaną w metodach ML, FIML itp.).
- $P(Y)$ często pomijamy jako stałą skalującą (gwarantującą że funkcja gęstości całkuje się do 1), zapisując proporcjonalność $P(\theta|Y) \propto P(Y|\theta) \cdot P(\theta)$.

Rozkład brzegowy *a posteriori*

Łączny rozkład *a posteriori* ma w modelu DSGE wiele wymiarów (tyle ile parametrów), na ogół więcej niż 2, co uniemożliwia jego prezentację i interpretację.

W praktyce interesują nas raczej rozkłady brzegowe ze względu na pojedynczy parametr θ_i , tzn. rozkład łączny przecałkowany po wszystkich parametrach z wyjątkiem jednego. Np. przy trzech parametrach $\theta_1, \theta_2, \theta_3$:

$$P(\theta_1|Y) \propto \iint P(Y|\theta) \cdot P(\theta) d\theta_2 d\theta_3$$

Nie możemy liczyć na istnienie analitycznych rozwiązań powyższego problemu – musimy numerycznie analizować rozkład łączny, przybliżając histogramy obrazujące rozkłady brzegowe *a posteriori*.

Metody bayesowskie – wyzwania

Estymacja bayesowska modelu DSGE wiąże się zatem z następującymi wyzwaniami dla ekonometrika:

- 1 wybór gęstości *a priori* $P(\theta)$ dla poszczególnych parametrów
- 2 konstrukcja funkcji wiarygodności $P(Y|\theta)$
- 3 określenie rozkładu *a posteriori*: łącznego i brzegowego dla pojedynczego parametru

Zajmijmy się tymi tematami po kolei.

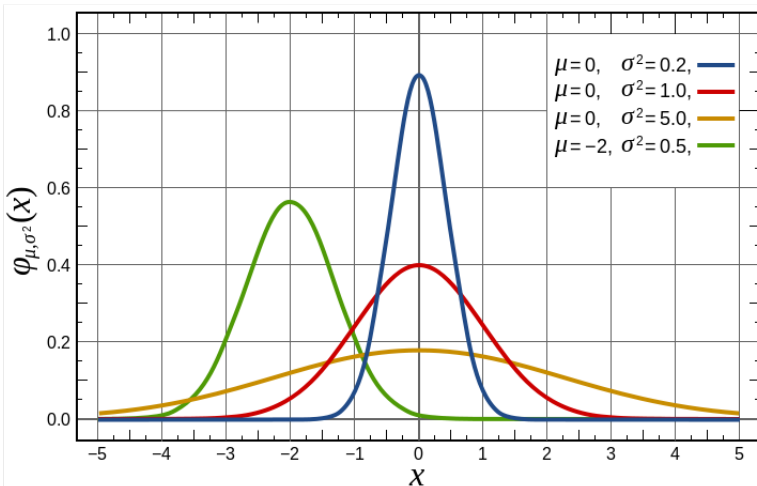
Plan prezentacji

- 1 Estymacja bayesowska
- 2 Rozkłady a priori
- 3 Konstrukcja funkcji wiarygodności
- 4 Rozkład a posteriori
- 5 W praktyce

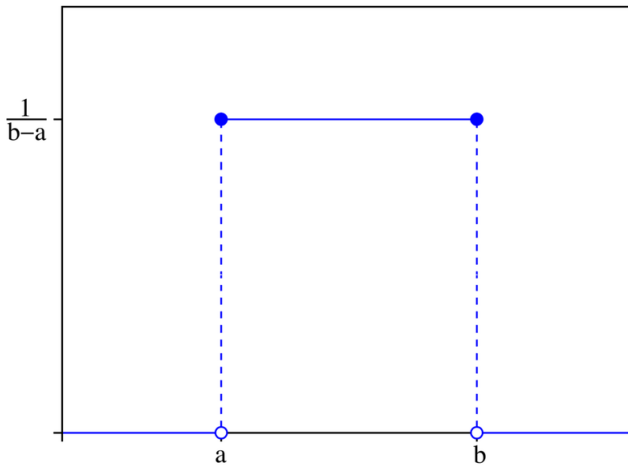
Typy rozkładów *a priori*

<u>PRIOR_SHAPE</u>	<u>Corresponding distribution</u>	<u>Range</u>
NORMAL_PDF	$N(\mu, \sigma)$	$\mathbb{R}$
GAMMA_PDF	$G_2(\mu, \sigma, p_3)$	$[p_3, +\infty)$
BETA_PDF	$B(\mu, \sigma, p_3, p_4)$	$[p_3, p_4]$
INV_GAMMA_PDF	$IG_1(\mu, \sigma)$	$\mathbb{R}^+$
UNIFORM_PDF	$U(p_3, p_4)$	$[p_3, p_4]$

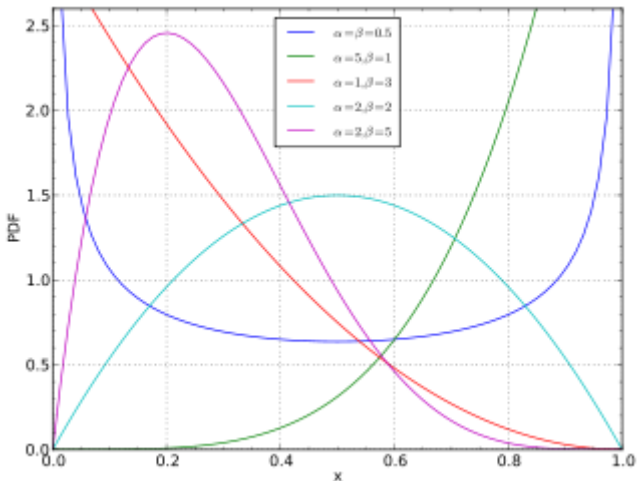
Rozkład normalny



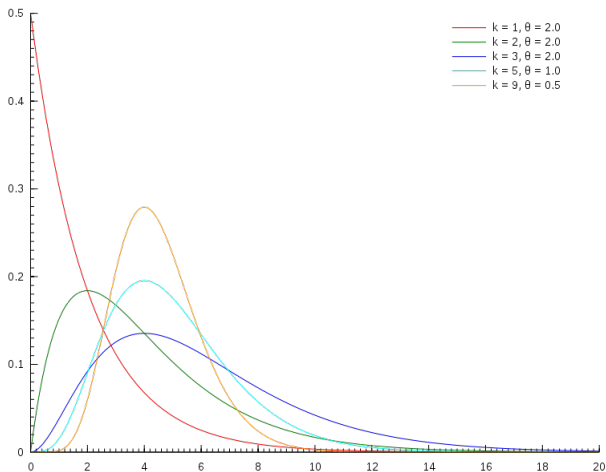
Rozkład jednostajny



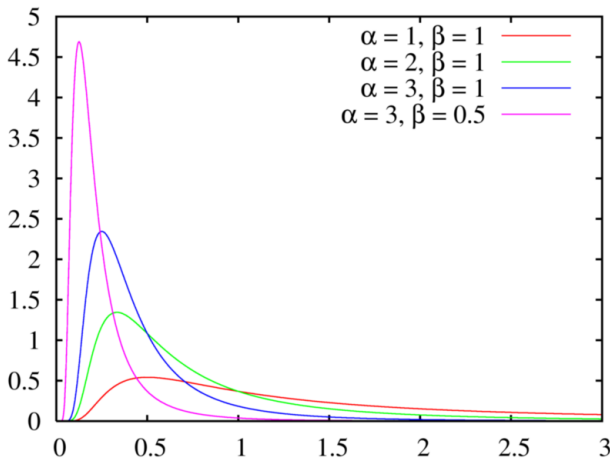
Rozkład beta



Rozkład gamma



Odwrotny rozkład gamma



Rozkłady *a priori* w Dynare

Blok **estimated_params** służy do określenia, które z parametrów są przedmiotem estymacji i określenia rozkładów *a priori*. Parametry, których tu nie ma, zostaną uznane za znane i skalibrowane w bloku **parameters**:

```
estimated_params;  
h, beta_pdf, 0.7, 0.1;  
sigma, gamma_pdf, 1.5, 0.4;  
correl_d, uniform_pdf, , , -0.99,0.99;  
sigma_d, inv_gamma_pdf, 4, inf;  
end;
```


Plan prezentacji

- 1 Estymacja bayesowska
- 2 Rozkłady a priori
- 3 Konstrukcja funkcji wiarygodności
- 4 Rozkład a posteriori
- 5 W praktyce

Filtr Kalmana – idea

- Zmienne w linearyzowanym modelu DSGE wchodzą do modelu w postaci odchyleń od stanu ustalonego. W takiej sytuacji jest bardzo prawdopodobne, że nie mierzymy ich dokładnie, bo nie znamy stanu ustalonego.
 - Co więcej, nawet PKB czy konsumpcja nie są mierzone stuprocentowo dokładnie (a jedynie szacowane) i błędy pomiaru mogą rzutować na dopasowanie naszej misternej sieci zależności do danych.
 - W danych uwzględnione są także zjawiska, których brak w naszym modelu. One nie zawsze pasują do naszych “strukturalnych” wstrząsów, które wbudowujemy w reguły decyzyjne podmiotów.

- Pojawia się zatem problem „ekstrakcji sygnału”, czyli identyfikacji i pomiaru zmiennej nieobserwowalnej na podstawie jej charakterystyki ekonomicznej, jak również sieci zależności między nią a obserwowalnymi.
 - Tę sieć zależności opisuje tzw. **model w przestrzeni stanów**.
- Przy estymacji pojawia się jednak trudność ze skonstruowaniem funkcji wiarygodności, bo występują w niej zmienne nieobserwowalne.
 - Aby przezwyciężyć problem, musimy posłużyć się tzw. **filtrem Kalmana**.

Model w przestrzeni stanów

Liniowy, Gaussowski:

$$\begin{cases} \alpha_t &= \mu + F\alpha_{t-1} + \varepsilon_t \\ y_t &= H\alpha_t + Ax_t + u_t \end{cases}$$

ZMIENNA STANU:

- α_t – wektor $m \times 1$ zmiennych modelu (nieobserwowalnych)
- $\varepsilon_t \sim N(0, Q)$ – wstrząsy w modelu

ZMIENNA SYGNAŁU:

- y_t – wektor $n \times 1$ zmiennych obserwowalnych
- $u_t \sim N(0, R)$ – błędy pomiaru

Rozkłady zmiennych stanu

$$\alpha_t = \mu + F\alpha_{t-1} + \varepsilon_t$$

- Przyjmujemy, że wektor α_{t-1} jest zmienną losową o wielowymiarowym rozkładzie normalnym (MVN) z wartością oczekiwaną $\mathbf{a}_{t-1}$ oraz wariancją $\mathbf{P}_{t-1}$.
- Wówczas α_t – jako kombinacja dwóch zmiennych o rozkładzie MVN – też ma rozkładm MVN.
- Pozostaje pytanie o rozkład α_1 (rekurencyjna definicja sprawia, że nie będziemy znali $\mathbf{a}_0$ ani $\mathbf{P}_0$).

Problem inicjalizacji

Przed okresem t znamy:

- y_{t-1}
- a_{t-1} – wartość oczekiwana α_{t-1} warunkowa względem y_{t-1}
- P_{t-1} – macierz wariancji-kowariancji α_{t-1} warunkowa względem y_{t-1}

Problem inicjalizacji: a_0 , P_0 – zwykle długookresowe wartości bezwarunkowe ([Hamilton, 1994](#)):

- $a_0 = \mu$
- $\text{vec}P_0 = [I_{m^2} - (F \otimes F)]^{-1} \cdot \text{vec}(Q)$

Możliwe również przyjęcie alternatywnego założenia: że α_0 ma tzw. rozkład nieinformacyjny (b. wysoka wariancja). Wówczas wyprowadzenie staje się nieco bardziej skomplikowane niż prezentowane tutaj.

Równania predykcji stanów

Co wiemy *ex ante* (przed t) na temat stanów w t ?

$$E_{t-1}(\alpha_t) = \mathbf{F} \cdot E_{t-1}(\alpha_{t-1}) + \mu = \mathbf{F}\mathbf{a}_{t-1} + \mu \equiv \mathbf{a}_{t|t-1}$$

$$\begin{aligned} E_{t-1} \left[(\alpha_t - \mathbf{a}_{t|t-1}) (\alpha_t - \mathbf{a}_{t|t-1})^T \right] &= \\ &= E_{t-1} \left[(\mathbf{F}\alpha_{t-1} + \mu + \varepsilon_t - \mathbf{F}\mathbf{a}_{t-1} - \mu) (\mathbf{F}\alpha_{t-1} + \mu + \varepsilon_t - \mathbf{F}\mathbf{a}_{t-1} - \mu)^T \right] = \\ &= E_{t-1} \left\{ [\mathbf{F}(\alpha_{t-1} - \mathbf{a}_{t-1}) + \varepsilon_t] [\mathbf{F}(\alpha_{t-1} - \mathbf{a}_{t-1}) + \varepsilon_t]^T \right\} = \\ &= \mathbf{F}\mathbf{P}_{t-1}\mathbf{F}^T + \mathbf{Q} \equiv \mathbf{P}_{t|t-1} \end{aligned}$$

Równania predykcji obserwacji

Co wiemy *ex ante* (przed t) na temat obserwacji w t ?

- $E_{t-1}(\mathbf{y}_t) = \mathbf{H}\mathbf{a}_{t|t-1} + \mathbf{A}\mathbf{x}_t \equiv \mathbf{y}_{t|t-1}$
- $E_{t-1}\left\{[\mathbf{y}_t - E_{t-1}(\mathbf{y}_t)][\mathbf{y}_t - E_{t-1}(\mathbf{y}_t)]^T\right\} = \mathbf{H}\mathbf{P}_{t|t-1}\mathbf{H}^T + \mathbf{R} \equiv \mathbf{V}_t$

Kowariancja *ex ante* stanu i obserwacji

$$\begin{aligned} E_{t-1} \left\{ [\alpha_t - \mathbf{a}_{t|t-1}] [\mathbf{y}_t - \mathbf{H}\mathbf{a}_{t|t-1} - \mathbf{A}\mathbf{x}_t]^T \right\} &= \\ &= E_{t-1} \left\{ [\alpha_t - \mathbf{a}_{t|t-1}] [\mathbf{H}\alpha_t + \mathbf{A}\mathbf{x}_t + \mathbf{u}_t - \mathbf{H}\mathbf{a}_{t|t-1} - \mathbf{A}\mathbf{x}_t]^T \right\} = \\ &= E_{t-1} \left\{ [\alpha_t - \mathbf{a}_{t|t-1}] [\alpha_t - \mathbf{a}_{t|t-1}]^T \mathbf{H}^T \right\} = \mathbf{P}_{t|t-1} \mathbf{H}^T \end{aligned}$$

Ta kowariancja uzupełnia opis łącznego rozkładu *ex ante* α_t i $\mathbf{y}_t$

jako $MVN \left(\begin{bmatrix} \mathbf{a}_{t|t-1} \\ \mathbf{y}_{t|t-1} \end{bmatrix}, \begin{bmatrix} \mathbf{P}_{t|t-1} & \mathbf{P}_{t|t-1} \mathbf{H}^T \\ \mathbf{H} \mathbf{P}_{t|t-1} & \mathbf{V}_t \end{bmatrix} \right).$

Równania aktualizacji

Czego się dowiedzieliśmy po zaobserwowaniu $\mathbf{y}_t$ (*ex post* w t):

- $\mathbf{e}_t = \mathbf{y}_t - \mathbf{y}_{t|t-1}$

Z perspektywy okresu t , $\mathbf{y}_t$ nie jest już zmienną losową, a więc rozkład α_t wyznaczamy na podstawie własności MVN jako rozkład warunkowy:

- $E_t(\alpha_t) = \mathbf{a}_{t|t-1} + \mathbf{P}_{t|t-1} \mathbf{H}^T \mathbf{V}_t^{-1} \mathbf{e}_t \equiv \mathbf{a}_t$

- $E_t \left[(\alpha_t - \mathbf{a}_t) (\alpha_t - \mathbf{a}_t)^T \right] =$
 $\mathbf{P}_{t|t-1} - \mathbf{P}_{t|t-1} \mathbf{H}^T \mathbf{V}_t^{-1} \mathbf{H} \mathbf{P}_{t|t-1} \equiv \mathbf{P}_t$

Funkcja wiarygodności

$$\mathcal{L}[y, \theta] = -\frac{Tn}{2} \ln(2\pi) + \frac{T}{2} \sum_{t=1}^T \ln |V_t(\theta)| - \frac{1}{2} \sum_{t=1}^T \mathbf{e}_t(\theta)^T V_t(\theta)^{-1} \mathbf{e}_t(\theta)$$
$$\rightarrow \max_{\theta}$$

Formy prezentacji zmiennych nieobserwowalnych

- Wartości zmiennych nieobserwowalnych (stanów) mogą być przybliżane punktowo poprzez $\hat{\mathbf{a}}_t$ (po oszacowaniu macierzy parametrów modelu)...
- ...lub przedziałowo, z dodatkowym wykorzystaniem $\hat{\mathbf{P}}_t$ (rozkład MVN dla α_t).
- Należy jednak pamiętać, że oba te wektory (tzw. zmienne **filtrowane**) bazują na wiedzy od okresu 1 do t .
- Ekonometryk dysponujący wiedzą do T włącznie jest w stanie dodatkowo wyznaczyć wartości **wygładzone** zmiennych nieobserwowalnych (wzory: zob. [Hamilton, 1994](#), roz. 13.6).

Model przestrzeni stanów i filtr Kalmana w Dynare

Dynare automatycznie określa model przestrzeni stanów i konstruuje funkcję wiarygodności, zakładając przy tym, że:

- ❶ wszystkie nasze zmienne w pliku .mod to zmienne nieobserwowalne (stany, α_t)
- ❷ mają one swoje “obserwowalne” odpowiedniki, a więc są mierzone w sposób niedoskonały
 - ❶ macierz Z jest jednostkowa
 - ❷ równania stanu to równania naszego modelu
- ❸ wstrząsy strukturalne oraz błędy pomiaru mają rozkład normalny

Zmienne obserwowalne w Dynare

Po bloku model w pliku `.mod` wpisujemy blok **varobs**, czyli które ze zmiennych endogenicznych są obserwowalne i będą wykorzystane w naszej analizie empirycznej:

```
varobs y i pi;
```

Dane muszą zostać dostarczone w pliku przestrzeni roboczej `.mat` (zob. wcześniejsze zajęcia). Plik ten musi zawierać pionowe wektory z danymi o jednakowej długości, nazywające się tak samo, jak nasze zmienne w bloku **varobs**.

Błędy pomiaru w Dynare

Jeżeli wprowadzamy błędy pomiaru w filtrze Kalmana, to dodatkowo w bloku **estimated_params** wpisujemy wiersze postaci:

```
stderr y, inv_gamma_pdf, 0.1, inf;  
stderr ir, inv_gamma_pdf, 0.1, inf;  
stderr pi, inv_gamma_pdf, 0.1, inf;
```

Plan prezentacji

- 1 Estymacja bayesowska
- 2 Rozkłady a priori
- 3 Konstrukcja funkcji wiarygodności
- 4 Rozkład a posteriori**
- 5 W praktyce

Algorytm Metropolisa-Hastingsa – przypadek ogólny

Rozważmy wektor parametrów $\theta = (\theta_1, \theta_2, \dots, \theta_K)$ o nieznaney gęstości a posteriori $p(\theta|\mathbf{y})$. Nie możemy również ustalić **rozkładów warunkowych**.

- 1 Wybieramy wektor startowy $\theta^{(0)} = (\theta_1, \theta_2, \dots, \theta_K)$.
- 2 Losujemy „kandydata” θ^* korzystając z gęstości generującej kandydatów, $q(\theta|\theta^{(0)})$.
- 3 Obliczamy prawdopodobieństwo akceptacji kandydata, $\alpha(\theta^*, \theta^{(0)})$.
- 4
$$\theta^{(1)} = \begin{cases} \theta^* & \text{z prawdopodobieństwem } \alpha(\theta^*, \theta^{(0)}) \\ \theta^{(0)} & \text{z prawdopodobieństwem } 1 - \alpha(\theta^*, \theta^{(0)}) \end{cases}$$
- 5 Powtarzamy kroki 1-4 z wektorem startowym $\theta^{(1)}$.
- 6 Powtarzamy sekwencję S razy.

Algorytm Metropolisa-Hastingsa – przypadek ogólny

Rozważmy wektor parametrów $\theta = (\theta_1, \theta_2, \dots, \theta_K)$ o nieznaney gęstości a posteriori $p(\theta|\mathbf{y})$. Nie możemy również ustalić **rozkładów warunkowych**.

- 1 Wybieramy wektor startowy $\theta^{(0)} = (\theta_1, \theta_2, \dots, \theta_K)$.
- 2 Losujemy „kandydata” θ^* korzystając z gęstości generującej kandydatów, $q(\theta|\theta^{(0)})$.
- 3 Obliczamy prawdopodobieństwo akceptacji kandydata, $\alpha(\theta^*, \theta^{(0)})$.
- 4
$$\theta^{(1)} = \begin{cases} \theta^* & \text{z prawdopodobieństwem } \alpha(\theta^*, \theta^{(0)}) \\ \theta^{(0)} & \text{z prawdopodobieństwem } 1 - \alpha(\theta^*, \theta^{(0)}) \end{cases}$$
- 5 Powtarzamy kroki 1-4 z wektorem startowym $\theta^{(1)}$.
- 6 Powtarzamy sekwencję S razy.

Algorytm Metropolisa-Hastingsa – przypadek ogólny

Rozważmy wektor parametrów $\theta = (\theta_1, \theta_2, \dots, \theta_K)$ o nieznaney gęstości a posteriori $p(\theta|\mathbf{y})$. Nie możemy również ustalić **rozkładów warunkowych**.

- 1 Wybieramy wektor startowy $\theta^{(0)} = (\theta_1, \theta_2, \dots, \theta_K)$.
- 2 Losujemy „kandydata” θ^* korzystając z gęstości generującej kandydatów, $q(\theta|\theta^{(0)})$.
- 3 Obliczamy prawdopodobieństwo akceptacji kandydata, $\alpha(\theta^*, \theta^{(0)})$.
- 4
$$\theta^{(1)} = \begin{cases} \theta^* & \text{z prawdopodobieństwem } \alpha(\theta^*, \theta^{(0)}) \\ \theta^{(0)} & \text{z prawdopodobieństwem } 1 - \alpha(\theta^*, \theta^{(0)}) \end{cases}$$
- 5 Powtarzamy kroki 1-4 z wektorem startowym $\theta^{(1)}$.
- 6 Powtarzamy sekwencję S razy.

Algorytm Metropolisa-Hastingsa – przypadek ogólny

Rozważmy wektor parametrów $\theta = (\theta_1, \theta_2, \dots, \theta_K)$ o nieznaney gęstości a posteriori $p(\theta|\mathbf{y})$. Nie możemy również ustalić **rozkładów warunkowych**.

- 1 Wybieramy wektor startowy $\theta^{(0)} = (\theta_1, \theta_2, \dots, \theta_K)$.
- 2 Losujemy „kandydata” θ^* korzystając z gęstości generującej kandydatów, $q(\theta|\theta^{(0)})$.
- 3 Obliczamy prawdopodobieństwo akceptacji kandydata, $\alpha(\theta^*, \theta^{(0)})$.
- 4
$$\theta^{(1)} = \begin{cases} \theta^* & \text{z prawdopodobieństwem } \alpha(\theta^*, \theta^{(0)}) \\ \theta^{(0)} & \text{z prawdopodobieństwem } 1 - \alpha(\theta^*, \theta^{(0)}) \end{cases}$$
- 5 Powtarzamy kroki 1-4 z wektorem startowym $\theta^{(1)}$.
- 6 Powtarzamy sekwencję S razy.

Algorytm Metropolisa-Hastingsa – przypadek ogólny

Rozważmy wektor parametrów $\theta = (\theta_1, \theta_2, \dots, \theta_K)$ o nieznaney gęstości a posteriori $p(\theta|\mathbf{y})$. Nie możemy również ustalić **rozkładów warunkowych**.

- 1 Wybieramy wektor startowy $\theta^{(0)} = (\theta_1, \theta_2, \dots, \theta_K)$.
- 2 Losujemy „kandydata” θ^* korzystając z gęstości generującej kandydatów, $q(\theta|\theta^{(0)})$.
- 3 Obliczamy prawdopodobieństwo akceptacji kandydata, $\alpha(\theta^*, \theta^{(0)})$.
- 4
$$\theta^{(1)} = \begin{cases} \theta^* & \text{z prawdopodobieństwem } \alpha(\theta^*, \theta^{(0)}) \\ \theta^{(0)} & \text{z prawdopodobieństwem } 1 - \alpha(\theta^*, \theta^{(0)}) \end{cases}$$
- 5 Powtarzamy kroki 1-4 z wektorem startowym $\theta^{(1)}$.
- 6 Powtarzamy sekwencję S razy.

Algorytm Metropolisa-Hastingsa – przypadek ogólny

Rozważmy wektor parametrów $\theta = (\theta_1, \theta_2, \dots, \theta_K)$ o nieznaney gęstości a posteriori $p(\theta|\mathbf{y})$. Nie możemy również ustalić **rozkładów warunkowych**.

- 1 Wybieramy wektor startowy $\theta^{(0)} = (\theta_1, \theta_2, \dots, \theta_K)$.
- 2 Losujemy „kandydata” θ^* korzystając z gęstości generującej kandydatów, $q(\theta|\theta^{(0)})$.
- 3 Obliczamy prawdopodobieństwo akceptacji kandydata, $\alpha(\theta^*, \theta^{(0)})$.
- 4
$$\theta^{(1)} = \begin{cases} \theta^* & \text{z prawdopodobieństwem } \alpha(\theta^*, \theta^{(0)}) \\ \theta^{(0)} & \text{z prawdopodobieństwem } 1 - \alpha(\theta^*, \theta^{(0)}) \end{cases}$$
- 5 Powtarzamy kroki 1-4 z wektorem startowym $\theta^{(1)}$.
- 6 Powtarzamy sekwencję S razy.

Algorytm MH – wybór wektora startowego

- ❶ Arbitralne wskazanie przez użytkownika.
- ❷ Losowanie z brzegowych rozkładów *a priori*, jeżeli możliwe.
- ❸ Wielokrotne wybory z analizą wrażliwości.
 - W każdym przypadku odcinamy pewną liczbę początkowych wartości jako *burn-in*. Przy prawidłowej zbieżności oraz odpowiednio wysokiej liczbie iteracji (S) nie powinno to mieć znaczenia.

Algorytm MH – gęstość generująca kandydatów

- 1 W ogólnym przypadku zakładamy, że jest zależna od bieżącego punktu $\theta^{(s)}$.
 - Najczęstszą implementacją jest **Random Walk MH**, gdzie losujemy $\epsilon \sim N(\theta^{(s)}, \Sigma)$ i rozważamy kandydata:

$$\theta^* = \theta^{(s)} + \epsilon$$

- 2 Nie musi tak jednak być. Wówczas mówimy o implementacji **Independence Chain MH** (nie omawiamy tutaj).

Algorytm MH – gęstość generująca kandydatów

- 1 W ogólnym przypadku zakładamy, że jest zależna od bieżącego punktu $\theta^{(s)}$.
 - Najczęstszą implementacją jest **Random Walk MH**, gdzie losujemy $\epsilon \sim N(\theta^{(s)}, \Sigma)$ i rozważamy kandydata:

$$\theta^* = \theta^{(s)} + \epsilon$$

- 2 Nie musi tak jednak być. Wówczas mówimy o implementacji **Independence Chain MH** (nie omawiamy tutaj).

Algorytm MH – prawdopodobieństwo akceptacji α

- ❶ Ponieważ q to funkcja umowna, korzystanie z niej bez dodatkowych korekt nie gwarantuje nam uzyskania sekwencji losowań przybliżających rozkład *a posteriori*.
 - ❶ Algorytm bez korekt zbyt często pozostaje w obszarach o wysokiej gęstości *a posteriori*.
 - ❷ W związku z tym musimy go skorygować, by dostatecznie dobrze „zwiedzić” całą dziedzinę parametrów.
- ❷ Korekta polega na nieakceptowaniu wszystkich kandydatów wylosowanych na podstawie gęstości q . W przypadku braku akceptacji, kolejnym elementem łańcucha jest kopia poprzedniego.
 - W implementacji **Random Walk** wzór na prawdopodobieństwo akceptacji zależy od gęstości *a posteriori* (p) dla wektorów: poprzedniego ($\theta^{(s)}$) oraz kandydata (θ^*). Szczegóły: Koop (roz. 5.5).

$$\alpha(\theta^*, \theta^{(0)}) = \min \left[\frac{p(\theta^* | y)}{p(\theta^{(s-1)} | y)}, 1 \right]$$

Algorytm MH – prawdopodobieństwo akceptacji α

- ❶ Ponieważ q to funkcja umowna, korzystanie z niej bez dodatkowych korekt nie gwarantuje nam uzyskania sekwencji losowań przybliżających rozkład *a posteriori*.
 - ❶ Algorytm bez korekt zbyt często pozostaje w obszarach o wysokiej gęstości *a posteriori*.
 - ❷ W związku z tym musimy go skorygować, by dostatecznie dobrze „zwiedzić” całą dziedzinę parametrów.
- ❷ Korekta polega na nieakceptowaniu wszystkich kandydatów wylosowanych na podstawie gęstości q . W przypadku braku akceptacji, kolejnym elementem łańcucha jest kopia poprzedniego.
 - W implementacji **Random Walk** wzór na prawdopodobieństwo akceptacji zależy od gęstości *a posteriori* (p) dla wektorów: poprzedniego ($\theta^{(s)}$) oraz kandydata (θ^*). Szczegóły: Koop (roz. 5.5).

$$\alpha(\theta^*, \theta^{(0)}) = \min \left[\frac{p(\theta^* | y)}{p(\theta^{(s-1)} | y)}, 1 \right]$$

Algorytm MH – α versus q

- Miarą jakości wyników jest m.in. średnie prawdopodobieństwo akceptacji $\bar{\alpha}$.
- Okazuje się, że optymalne wartości $\bar{\alpha} \in [0, 2; 0, 4]$. Dostatecznie niskie prawdopodobieństwo akceptacji oznacza, że dziedzina gęstości *a posteriori* została dobrze wyeksplorowana.
- $\bar{\alpha}$ to jednak wartość wynikowa i nie możemy jej wprost wybrać. Zależy ona przede wszystkim od doboru gęstości generującej kandydatów q .
 - W przypadku **Random Walk MH**, sprowadza się to do odpowiedniego ustalenia wariancji kroku ε , czyli Σ (w Dynare: **mh_jscale**)
 - Relację między $\bar{\alpha}$ a Σ należy zbadać w ramach dodatkowej procedury iteracyjnej.
 - Zaczynamy w niej od $\Sigma^{(0)} = c^{(0)} \cdot I$. W przypadku zbyt wysokiego $\bar{\alpha}^{(0)}$ zbyt często akceptujemy, a więc jesteśmy zbyt konserwatywni w „zwiedzaniu” dziedziny, czyli powinniśmy ustalić $c^{(1)} > c^{(0)}$.

Algorytm MH w Dynare (1)

Polecenie estymacji jest konstruowane w następujący sposób:

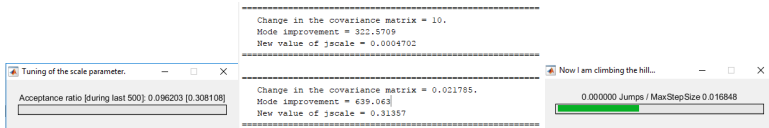
```
estimation(datafile = plik_danych, mh_replic = 20000,  
mh_nblocks = 5, mh_drop = 0.5, mh_jscale = 0.2);
```

Kluczowe opcje:

- **plik_danych.mod** to plik w przestrzeni roboczej zawierający wektory zmiennych obserwowalnych;
- **mh_replic** – liczba replikacji algorytmu (długość łańcucha);
- **mh_nblocks** – liczba łańcuchów (co najmniej 2 – ale lepiej więcej, np. 5-10, tyle że to jest czasochłonne...);
- **mh_drop** – frakcja odrzuconych fragmentów łańcucha blisko początku (by uniezależnić się od losowo wybranego początku).
- dalsze opcje w manualu Dynare (np. dobór próby)

Algorytm MH w Dynare (2)

`mh_jscale` to parametr powiązany z wyborem odpowiedniej wariancji kroku i podlega ew. korekcie w przebiegu estymacji (następnym razem możemy podać inną wartość, aby przyspieszyć):



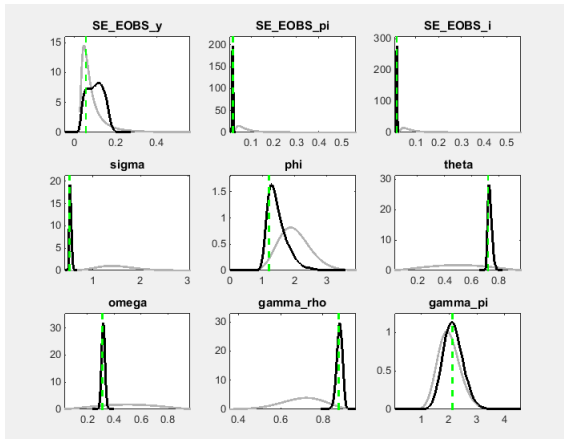
Jak ocenić wyniki estymacji? (1a)

- Rolę oszacowania punktowego pełni na ogół średnia z rozkładu *a posteriori*. Możemy zastanowić się nad jej wymiarem ekonomicznym i ulokowaniem na tle naszej wiedzy *a priori*.
- Statystyczny “sukces” osiągamy, gdy rozkład *a posteriori* widocznie różni się od rozkładu *a priori* bardziej smukłym kształtem (niższą wariancją) – oznacza to, że próba sprecyzowała naszą wiedzę o parametrze w stosunku do tej, którą posiadaliśmy wcześniej. Rozkłady *a posteriori* powinny mieć ponadto stosunkowo regularny kształt (np. być jednomodalne).

Jak ocenić wyniki estymacji? (1a)

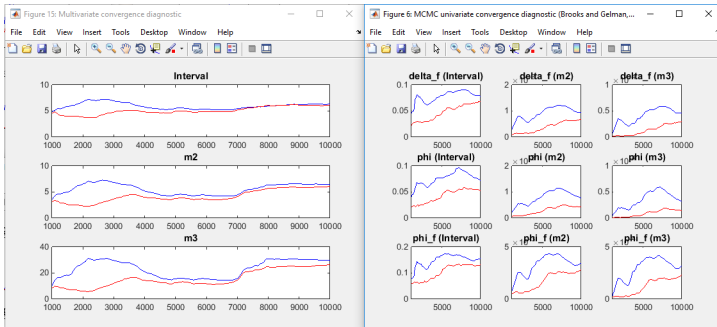
- Rolę oszacowania punkowego pełni na ogół średnia z rozkładu *a posteriori*. Możemy zastanowić się nad jej wymiarem ekonomicznym i ulokowaniem na tle naszej wiedzy *a priori*.
- Statystyczny “sukces” osiągamy, gdy rozkład *a posteriori* widocznie różni się od rozkładu *a priori* bardziej smukłym kształtem (niższą wariancją) – oznacza to, że próba sprecyzowała naszą wiedzę o parametrze w stosunku do tej, którą posiadaliśmy wcześniej. Rozkłady *a posteriori* powinny mieć ponadto stosunkowo regularny kształt (np. być jednomodalne).

Jak ocenić wyniki estymacji? (1b)



Jak ocenić wyniki estymacji? (2a)

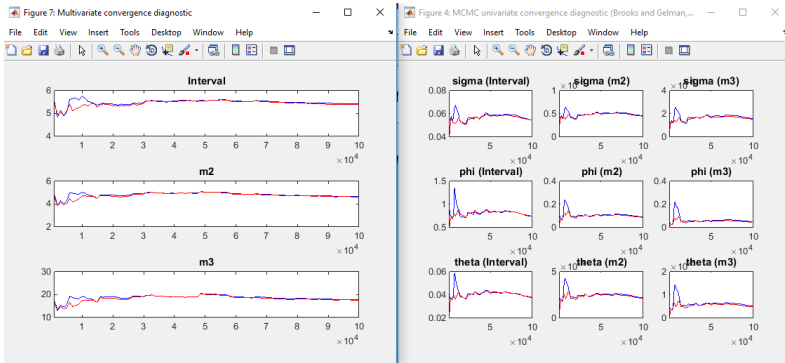
Czerwona linia: wariancja wewnątrzłańcuchowa (**within**). Niebieska linia: całkowita wariancja (**within + between**).



Jak ocenić wyniki estymacji? (2b)

Stabilizacja **czzerwonej** linii: losowanie z tego samego rozkładu (symptom osiągnięcia zbieżności).

Zbieżność **obu** linii: brak wariancji międzyłańcuchowej (symptom osiągnięcia zbieżności).

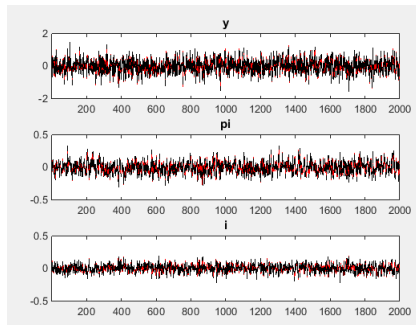


Jak ocenić wyniki estymacji? (3a)

- Obejrzyjmy również wnikliwie ścieżki zmiennych obserwowalnych (nasze dane) vs wygładzonych (tzn. pozbawionych błędów pomiaru).
 - Gdy zmienne wygładzone różnią się bardzo, np. usunięta jest z nich większość wariancji, mogliśmy osiągnąć efekt odwrotny do zamierzonego.
 - Wariancja zawarta w tych zmiennych nie pasowała do modelu, a więc została ona “wypchnięta” poza niego i wykorzystane dane w zasadzie w ogóle nie wzięły udziału w ustalaniu wartości parametrów strukturalnych.
 - Innymi słowy, model obraził się na rzeczywistość, że do niego nie przystaje, i “się w sobie zamknął”... (tzn. przypisał lwia część wariancji zmiennych błędom pomiaru – empiryczna porażka modelu!).

Jak ocenić wyniki estymacji? (3b)

Czerwona linia (bez błędu pomiaru) powinna mieć gładniejszy przebieg niż czarna (z błędem pomiaru). Jednak nie powinna być zupełnie płaska w porównaniu z czarną, bo to sugeruje kłopot z dopasowaniem.



Plan prezentacji

- 1 Estymacja bayesowska
- 2 Rozkłady a priori
- 3 Konstrukcja funkcji wiarygodności
- 4 Rozkład a posteriori
- 5 W praktyce

Zmienne w modelu DSGE (1)

Zanim przystąpimy do estymacji, przypomnijmy sobie, że...

- wiele zmiennych w modelu (y , c) to procentowe odchylenia od stanu ustalonego (gdzie ich szukać?)
- inflacja (π) mierzona jest wokół stanu ustalonego w którym ceny nie zmieniają się (a cel inflacyjny NBP?)
- produkcja i konsumpcja (y , c) mierzą ilości, a więc interesują nas kategorie realne (*volumes*), a nie nominalne
- stopy procentowe (i) winny być wyrażone w rzędzie wielkości analogicznym do procentowych odchylenia od stanu ustalonego, tzn. konsekwentnie 1 albo 0.01 dla stopy procentowej równej 1% i jednoprocentowego odchylenia od stanu ustalonego dla produkcji
- nie ma czegoś takiego jak dane o sektorze T i NT, musimy je sobie sami skonstruować...
- w modelu nie mamy wbudowanych żadnych mechanizmów sezonowości (np. egzogeniczna pogoda czy schematy zatrudnienia), dlatego korzystamy z danych odsezonowanych przez Eurostat

Zmienne w modelu DSGE (2)

A zatem będziemy musieli wykonać kilka operacji:

- ściągnięcie odpowiednich szeregów (zob. **dane.xls**, zakładka **Eurostat**)
- kalibracja parametrów, których nie będziemy estymować (zob. **dane.xls**, 2. zakładka – kolumny M i N, 3. zakładka – kolumny B, C i D, 4. zakładka)
- konstrukcja zmiennych sektorowych (y , π) na podstawie danych o wartości dodanej i jej deflatorze (**dane.xls**, zakładka **dane_do_matlaba** i poprzednie)
- logarytmowanie indeksów ilości (y , c) oraz cen i kursu walutowego, wyrażenie tych ostatnich jako logarytmicznych przyrostów (zakładka **dane_do_matlaba**)
- odjęcie średniej od wszystkich szeregów (zob. **import_danych_matlab.m**)

Zmienne w modelu DSGE (3)

- wyrażenie y oraz c jako odchyłeń od stanu ustalonego poprzez logarytmiczne odchylenia od trendu wyznaczonego filtrem Hodricka-Prescotta
 - w Matlabie jest polecenie `hpfilter` w **Econometrics Toolbox**
 - jeżeli nie mamy tego dodatku lub mamy Octave, korzystamy z dołączonego pliku **hpfilter.m**
 - filtrację można też przeprowadzić łatwo w innym programie, np. w Gretlu
- w naszym przykładzie filtrujemy również stopę nominalną w Polsce ze względu na obniżający się cel inflacyjny w początkowym okresie próby

Ćwiczenie 1

Zanim przejdziemy do prawdziwej konfrontacji z danymi (która może rozczarować), bezpiecznym ćwiczeniem będzie:

- 1 symulacja stochastyczna modelu gospodarki zamkniętej;
- 2 zapisanie wygenerowanych ścieżek;
- 3 estymacja parametrów modelu na podstawie tych ścieżek.

Całość zapiszemy jednym pliku modelu (mod).

Uwaga! Potrzebne będzie polecenie zapisujące wygenerowane przez Dynare zmienne do pliku zewnętrznego Octave/Matlab (jako wektory):

```
dynasave('plik_danych.mat') y i pi;
```

Ćwiczenie 2

Przeprowadź estymację modelu z pliku **NK_open_estimation.mod** dla strefy euro oraz kraju innego niż Polska, np. Czech. Samodzielnie skonstruuuj zbiór danych i dokonaj odpowiednich korekt.

Ćwiczenie 3

Przeprowadź estymację dla modelu gospodarki otwartej należącego do strefy euro od początku, np. Irlandii. Przyjmij w uproszczeniu, że dane dla całej strefy euro nie obejmują tego kraju (lub, innymi słowy, że to kraj na tyle mały, by nie ważyć na nich istotnie). Jakich zmian należy dokonać w modelu i zbiorze danych? (przypomnijmy sobie wynik jednego z zadań z poprzedniego tematu)