

MATERIAŁY I STUDIA

Zeszyt nr 273

Wykresy wachlarzowe inflacji
a różne wymiary niepewności

Halina Kowalczyk

Warszawa, 2012 r.

Dziękuję kolegom z Instytutu Ekonomicznego NBP za inspirację do podjęcia tematu,
a Tomaszowi Łyziakowi i Recenzentowi za cenne uwagi i sugestie.

Poglądy przedstawione w artykule są poglądami własnymi autorki i nie muszą odzwierciedlać stanowiska NBP.

Projekt graficzny:
Oliwka s.c.

Skład i druk:
Drukarnia NBP

Wydął:
Narodowy Bank Polski
Departament Edukacji i Wydawnictw
00-919 Warszawa, ul. Świętokrzyska 11/21
tel. 22 653 23 35, fax 22 653 13 21

© Copyright Narodowy Bank Polski, 2012

ISSN 2084–6258

Materiały i Studia są rozprowadzane bezpłatnie

Dostępne są również na stronie internetowej NBP: <http://www.nbp.pl>

Spis treści:

1. Zapomniane funkcje wykresów wachlarzowych.....	3
Pierwotne znaczenie pojęcia „fan chart”	4
Wykresy wachlarzowe tworzone w oparciu o błędy statystyczne.....	5
2. Systematyka niepewności oparta na trzech wymiarach.....	9
3. Drugi wymiar niepewności w popularnych metodach.....	12
4. Trzeci wymiar niepewności a dotychczasowe metody.....	14
5. Problem wymieszania różnych typów niepewności.....	18
6. Propozycja innego podejścia do konstrukcji wykresów wachlarzowych.....	20
Dyskretny rozkład subiektywny.....	21
Ciągły rozkład subiektywny.....	22
Szczególny przypadek ciągłego rozkładu subiektywnego – TPN.....	23
Przykład wykresu wachlarzowego – dwa alternatywne scenariusze.....	27
7. Zakończenie.....	30
 Bibliografia	 31

Streszczenie

W opracowaniu omówiono problemy związane z przedstawianiem niepewności prognoz makroekonomicznych przy pomocy wykresów wachlarzowych (*fan charts*). Zwrócono uwagę na konieczność odróżniania błędów statystycznych i zmienności w czasie od niepewności wynikającej z niepełnej wiedzy. Wskazano na potrzebę dokładniejszego definiowania składowych niepewności. Pokazano ograniczenia tradycyjnych metod konstrukcji wykresów wachlarzowych w sytuacjach, gdy istnieje duża niepewność co do możliwych scenariuszy.

Zaproponowano modyfikacje polegające na wprowadzeniu dwóch rozkładów, z których jeden opisuje niepewność scenariuszową, a drugi charakteryzuje niepewność modelu lub systemu prognostycznego. Do przedstawiania całkowitej niepewności wykorzystywane są mieszanki rozkładów albo sploty gęstości prawdopodobieństwa. Proponowane podejście pozwala na zachowanie informacji o scenariuszach istotnych dla procesu decyzyjnego, a także na osobne analizowanie niepewności o różnym charakterze.

Słowa kluczowe: wykresy wachlarzowe, projekcje inflacji, subiektywne rozkłady prawdopodobieństwa, niepewność, scenariusze, splot gęstości

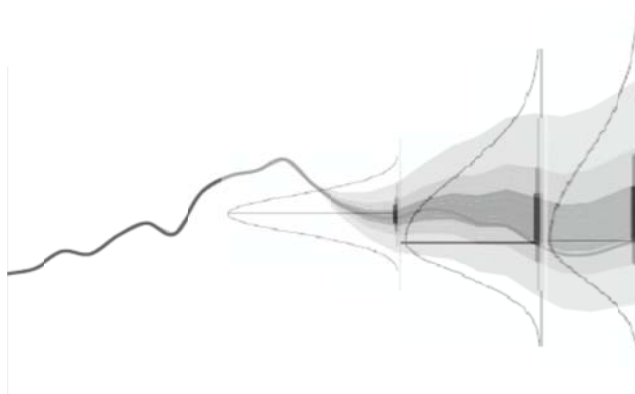
Klasyfikacja JEL: C19, C53, E39

1. Zapomniane funkcje wykresów wachlarzowych

Procesowi prognostycznemu towarzyszy niepewność – wykorzystujemy modele, które są jedynie uproszczonym obrazem rzeczywistości; nie dysponujemy dokładnymi danymi; nie wiemy, co się wydarzy. Od lat powszechnie wykorzystywanym sposobem prezentowania niepewności projekcji głównych wskaźników makroekonomicznych, publikowanych przez banki centralne, są wykresy wachlarzowe (*fan charts*). NBP idąc za przykładem Banku Anglii i Banku Szwecji, umieszcza je w „Raportach o Inflacji” od 2004 roku. Wykresy wachlarzowe konstruowane są w oparciu o teoretyczne lub empiryczne funkcje gęstości rozkładów prawdopodobieństwa, otrzymane w wyniku analizy niepewności dla przedstawianych horyzontów prognozy (wykres 1).

Wykres 1

Wykres wachlarzowy a funkcje gęstości



Sposób uzyskiwania rozkładów, tzn. wiedza o tym jakie przyjęto założenia i jaką niepewność uwzględniono w obliczeniach, ma zasadnicze znaczenie dla właściwej interpretacji wykresów wachlarzowych. Stało się to szczególnie ważne w kontekście poszerzenia spektrum modeli wykorzystywanych w bankach, bowiem zakres możliwej do przeprowadzenia analizy niepewności jest w znacznej mierze zdeterminowany typem modelu. Nie bez znaczenia są też preferencje zmieniających się zespołów prognostycznych.

Skupiając uwagę na samym wykresie trudno jest na ogół stwierdzić, czy przedstawiane są przedziały ufności otrzymane w wyniku estymacji, przewidywana zmienność, czy też możliwość wystąpienia różnych wartości wskaźnika odpowiadających odmiennym scenariuszom rozwoju sytuacji gospodarczej.

Pierwotne znaczenie pojęcia „fan chart”

Koncepcja wykresów wachlarzowych narodziła się w Banku Anglii. Pierwszy *fan chart* został przedstawiony w *Inflation Report* w lutym 1996 roku, ale już wcześniej Bank Anglii uznawał za konieczne zwrócenie uwagi na towarzyszącą prognozom niepewność i w latach 1993-1995 publikował swoje projekcje w formie wykresów przedstawiających ścieżkę centralną otoczoną zacieniowanym (jednokolorowym) obszarem o granicach wyznaczonych przez średni błąd absolutny prognoz z poprzednich 10 lat.

Ten sposób prezentacji uznano jednak szybko za niesatysfakcjonujący (por. Britton, Fisher, Whitley 1998) z następujących powodów:

Odbiorcy prognoz ignorowali informację o możliwym dużym błędzie i nadal przywiązywali nadmierną wagę do ścieżki centralnej. Granice przedstawianego obszaru niepewności były często błędnie interpretowane jako zakres wartości prognozowanego wskaźnika. Brakowało, istotnej dla uzasadnienia podejmowanych przez Bank decyzji, oceny ryzyka realizacji scenariuszy odmiennych niż ten, któremu odpowiadała publikowana ścieżka centralna.

Wprowadzenie wykresów wachlarzowych stworzyło możliwość odzwierciedlenia subiektywnej oceny Banku Anglii na temat presji inflacyjnej i jej rozwoju w czasie. Błąd statystyczny zastąpiono rozkładem prawdopodobieństwa odzwierciedlającym dyskusję nad alternatywnymi założeniami i pokazującym, jak prawdopodobne są wg członków *Monetary Policy Committee (MPC)* różne wartości wskaźnika inflacji (przy niezmienionych stopach procentowych). Wykresy wachlarzowe pokazywały prawdopodobieństwo przedziałów zbudowanych wokół prognozy centralnej otrzymanej w wyniku przeliczenia modelu dla najbardziej prawdopodobnego scenariusza oraz uwzględniały inne - relatywnie mniej prawdopodobne scenariusze prowadzące do ścieżek o prawdopodobieństwie wynikającym z subiektywnych ocen dotyczących szans realizacji odpowiedniego zestawu założeń. Otrzymanie pełnego rozkładu wymagałoby przeliczenia modelu dla wszystkich alternatywnych założeń, co oczywiście było nierealne. W praktyce analizowano ich ograniczoną liczbę i dopasowywano do wyników rozkład prawdopodobieństwa.

O problemach wynikających z postaci zakładanego rozkładu będzie mowa w rozdz. 4. W tym miejscu chcemy jedynie podkreślić, że pierwotnie pod pojęciem *fan chart* rozumiano nie tylko typ wykresu, tzn. sposób graficznego przedstawienia rozkładu prawdopodobieństwa. Co najmniej równie istotne były przypisane wykresowi funkcje - nie sprowadzały się one do statystycznego opisu błędów i zmienności. Brano je pod uwagę w trakcie określania

parametrów rozkładów, ale bardzo ważne było odzwierciedlenie niepewności wynikającej z ograniczonych możliwości przewidywania przyszłych zjawisk makroekonomicznych.

O tej znaczeniowej warstwie nie zawsze się pamięta.

Wykresy wachlarzowe tworzone w oparciu o błędy statystyczne

Z czasem wykresy wachlarzowe zaczęto traktować jako uniwersalne narzędzie do przedstawiania prognoz makroekonomicznych otrzymywanych przy pomocy różnorodnych modeli pełniących zarówno główne, jak i pomocnicze funkcje w systemach prognostycznych. Są wśród nich modele, które nie odwołują się do przyszłych wartości zmiennych egzogenicznych, a więc nie pozwalają na bezpośrednie przeprowadzanie analiz scenariuszowych stanowiących istotę wykresów wachlarzowych w ich pierwotnym wydaniu. Do tego typu modeli należy wiele prostych modeli statystycznych, ale też złożone modele DSGE, w których jedynymi zmiennymi określanymi poza modelem są losowe szoki.

Być może właśnie rosnącą rolą modeli DSGE w systemach prognostycznych banków centralnych należy tłumaczyć pewien regres w zakresie komunikowania niepewności przejawiający się w powrocie do charakteryzowania niepewności prognoz za pomocą błędów statystycznych *ex post*. Dotyczy to nawet oficjalnej projekcji Banku Szwecji, w którym wcześniej przywiązywano dużą wagę do sygnalizowania ryzyka wystąpienia odmiennych scenariuszy, a od 2007 r. zaczęto publikować idealnie symetryczne wykresy wachlarzowe oparte na błędach historycznych. Sytuacja jest o tyle paradoksalna, że takie tendencje obserwujemy w okresie upowszechniania się metod otwierających nowe możliwości, w szczególności metod wykorzystujących podejście bayesowskie czy filtr Kalmana.

Historyczne błędy statystyczne, w tym najpopularniejszy RMSE (pierwiastek błędu średniokwadratowego) są oczywiście dobrym punktem wyjścia do szacowania niepewności, gdyż uwzględniają różne źródła błędów popełnianych w przeszłości: niepewność modelu (dotyczącą zarówno struktury jak i oszacowań parametrów), błędy założeń, błędy pomiaru wielkości wykorzystywanych przez modele. Ale wykres wachlarzowy powstały w oparciu o RMSE informuje jedynie o tym, że twórcy prognoz mają świadomość dotychczasowych błędów towarzyszących ich prognozom punktowym i że analogiczne błędy mogą dotyczyć aktualnie przedstawianej projekcji centralnej. Taka informacja jest ważna dla uświadomienia odbiorcom prognoz, że są one z natury niepewne i konieczne jest uwzględnienie ryzyka łączącego się z ich wykorzystywaniem. Daje też możliwość oszacowania tego ryzyka i pozwala wnioskować o cechach modelu lub systemu prognostycznego. Jest natomiast niewy-

starczająca z punktu widzenia osób podejmujących decyzje, szczególnie wtedy, gdy w otoczeniu gospodarczym mają miejsce nietypowe zjawiska, gdyż informuje jedynie o przeciętnej wartości przeszłych błędów a nie o błędzie w konkretnym punkcie czasowym. Mówi o tym, jakiej dokładności możemy się spodziewać prognozując dotychczasowymi metodami i w zbliżonym środowisku. Nie ma jednak żadnej gwarancji, że przyszłe prognozy będą obarczone analogicznym błędem. RMSE nie dostarcza ponadto żadnej informacji o strukturze błędów. Będzie ona zupełnie różna dla różnych klas modeli.

Wykresy wachlarzowe oparte na RMSE, mogą być wykorzystywane do przedstawienia własności prognostycznych modeli w różnych horyzontach i być pomocne przy podejmowaniu decyzji o wyborze tego czy innego modelu dla określonych celów. Mogą stanowić podstawę do porównań modeli (o ile modele nie różnią się zbyt mocno złożonością czy stopniem egzogeniczności), gdyż jest dosyć obiektywną i uniwersalną miarą niepewności prognoz (oczywiście jedynie punktowych), co staje się dosyć istotne wobec różnorodności zarówno samych modeli jak też technik estymacji, filtracji i predykcji stosowanych w prognozowaniu makroekonomicznym. Jeżeli jednak decyzje nie dotyczą wyboru modelu, lecz sposobu prowadzenia polityki pieniężnej, to ich przydatność jest ograniczona, gdyż nie dają odpowiedzi na pytania o prawdopodobieństwo znalezienia się wskaźnika makroekonomicznego w tym czy innym przedziale dla określonego horyzontu prognozy.

Kolejnym problemem jest to, że RMSE nie odzwierciedla propagacji błędów w czasie. Otrzymujemy obraz błędów bezwarunkowych. Może się zdarzyć, szczególnie w przypadku modeli do prognozowania średnio- i długookresowego, że otrzymamy zwiężający się wachlarz ze względu na mniejsze błędy RMSE dla dłuższych horyzontów. (Charemza i in. (2009) pokazują jak inaczej można opisywać statystyczne błędy prognoz uwzględniając zależność i asymetrię rozkładów).

W przypadku pewnych klas modeli za naturalne można by uznać tworzenie wykresów wachlarzowych na podstawie teoretycznych rozkładów predyktywnych wyprowadzanych na podstawie postaci modelu. (Przegląd metod konstrukcji przedziałów predykcji dla wybranych typów modeli statystycznych można znaleźć np. w Chatfield 1993). Wiążą się z tym jednak pewne problemy, mianowicie:

Przedziały predykcji, otrzymywane w sposób klasyczny, bazują na dopasowaniu modelu do danych i opisują jedynie efekty niedokładności estymacji parametrów i/lub nie-

pewność rezydualną. Dla przykładu, przy wyprowadzaniu rozkładów predyktywnych dla modeli regresji liniowej, standardem stało się uwzględnianie niepewności oszacowań parametrów i niepewności rezydualnej, ale już w przypadku modeli typu ARIMA uwzględnia się zazwyczaj jedynie niepewność rezydualną - niepewność parametrów jest ignorowana, mimo znanych rezultatów teoretycznych (np. Fuller, Hasza 1981; de Luna 2000).

Rozkłady predyktywne nie uwzględniają niepewności zmiennych egzogenicznych. Jest to zresztą zrozumiałe wobec komplikacji, które byłyby z tym związane nawet w przypadku bardzo prostych modeli (por. Feldstein 1997).

Rozkładowi predyktywnym towarzyszy założenie o poprawności postaci modelu. Waga tego założenia jest różna w zależności od stopnia weryfikacji teorii leżącej u podstaw modelu. W przypadku modeli czysto statystycznych niepewność postaci modelu może być składową o zasadniczym znaczeniu (por. Chatfield 1995). Remedium może stanowić bayesowskie uśrednianie modeli (*Bayesian Model Averaging, BMA*; np. Hodges 1987; Raftery, Madigan, Volinsky, 1995). BMA jest z powodzeniem stosowane do wielu typów modeli statystycznych. Nie prowadzi jednak do interpretowalnego modelu, co ogranicza obszar zastosowań.

Zatem konstruowanie wykresów wachlarzowych w oparciu o teoretyczne rozkłady predyktywne nie jest raczej właściwym wyborem w kontekście decyzyjnym (nie tylko z tego powodu, że modele, dla których takie rozkłady można otrzymać, są zbyt uproszczone i mogą być wykorzystywane jedynie w charakterze pomocniczych).

Od modeli, których zadaniem jest wspieranie procesu decyzyjnego wymaga się znacznie więcej niż dobrego dopasowania do danych. Oczekuje się, że będą objaśniały zarówno to, co się już wydarzyło, jak i to co się może wydarzyć. Powinny też dawać możliwość analizowania różnorodnych implikacji obserwowanych lub przewidywanych zdarzeń oraz skutków odmiennych decyzji. Zespoły odpowiedzialne za przygotowanie projekcji na potrzeby polityki pieniężnej stają przed problemem uwzględnienia istotnych źródeł niepewności związanych z formułowanymi przez ekspertów założeniami dotyczącymi stanu początkowego (tzw. punktu startowego) i przyszłych wartości zmiennych opisujących otoczenie zewnętrzne. Pojawia się potrzeba włączenia niepewności o odmiennym charakterze – niepewności, która wynika z niepełnej wiedzy i może być opisana jedynie przy pomocy subiektywnych rozkładów prawdopodobieństwa, które nie odzwierciedlają reguł rządzących zjawiskami losowymi (w przeciwieństwie do obiektywnych rozkładów statystycznych).

Reasumując:

Rosnąca popularność wykresów wachlarzowych – traktowanie ich jako uniwersalnego narzędzia prezentowania prognoz makroekonomicznych – wymaga dokładniejszego określania, co kryje się pod pojęciem „niepewność”. Różne zespoły progностyczne mogą definiować niepewność w całkowicie odmienny sposób.

Projekcjom publikowanym przez banki centralne towarzyszą zwykle wyczerpujące opisy modeli i sposobu otrzymywania wykresów wachlarzowych (NBP publikuje szczegółowe informacje na swojej stronie internetowej). Zapoznanie się z objaśnieniami jest niezbędne dla właściwej interpretacji prognoz przedstawianych w kategoriach probabilistycznych. Ma też zasadnicze znaczenie dla wykorzystania informacji o niepewności przy podejmowaniu decyzji przez różne podmioty.

2. Systematyka niepewności oparta na trzech wymiarach

Wielu problemów, występujących przy analizowaniu i opisywaniu niepewności, można byłoby prawdopodobnie uniknąć uwzględniając jej wielowymiarowy charakter, jak to proponują Walker i in. (2003). Stworzyli oni typologię niepewności obejmującą, rozwijającą i porządkującą wiele innych systemów pojęciowych opracowanych przez specjalistów z różnych dziedzin. Wydaje się, że warto ją przenieść na grunt prognozowania makroekonomicznego.

Walker i in. pojmują niepewność bardzo szeroko - jest nią każde odchylenie od nieosiągalnego ideału, jakim jest determinizm i rozpatrują jej trzy wymiary. Pierwszy wymiar służy określeniu obszaru modelowania, w którym zlokalizowana jest analizowana składowa niepewności. Dwa kolejne określają jej poziom i naturę.

Lokalizacja niepewności

Wyróżniane są następujące lokalizacje:

Kontekst – określa, co jest przedmiotem modelowania, jaka część realnego świata jest opisywana przez model i jak kompletny jest to opis. Kontekst decyduje o możliwościach wykorzystania modelu do określonych celów. Użycie modelu w niewłaściwym kontekście wiąże się z niepewnością wyników.

Model – jego struktura (definicje zmiennych, relacje między nimi, założenia, algorytmy matematyczne) a także parametry i implementacja komputerowa.

Obszar wejścia – dane opisujące modelowany system i zmienne zewnętrzne wywołujące zmiany w systemie z podziałem na te, na które nie można wpływać i te, które podlegają kontroli czy sterowaniu (w oryginale: *scenario* i *policy variables*).

Obszar wyjścia - niepewność związana z tym obszarem, to efekt kumulacji niepewności z poprzednich obszarów. Poszczególne składowe propagowane są zgodnie z logiką modelu, dając w efekcie niepewność wyników modelu a więc to, co zwykle nazywane jest błędem predykcji.

Poziom niepewności

Proponowana jest następująca gradacja poziomu niepewności:

determinizm → niepewność statystyczna → niepewność scenariuszowa → uświadomiona niewiedza → absolutna niewiedza.

Determinizm odpowiada idealnej sytuacji, tzn. pewności. Na drugim krańcu jest sytuacja, gdy nie jesteśmy w stanie określić nawet tego, czego nie wiemy.

Niepewność statystyczna to poziom, przy którym możliwe jest scharakteryzowanie odchylenia od prawdziwej wartości przy pomocy formuł statystycznych. Przyjęcie tego poziomu niepewności związane jest z założeniem, że model dosyć wiernie opisuje analizowane zjawisko, oraz że dane użyte do estymacji/kalibracji modelu są reprezentatywne dla okoliczności, w których model będzie zastosowany.

Niepewność scenariuszowa jest poziomem, przy którym mamy do czynienia z możliwością wystąpienia różnych wartości, lecz mechanizm prowadzący do tych wartości nie jest na tyle rozpoznany, by mógł być opisany statystycznie. Scenariusze mówią nie o tym, co się zdarzy, lecz co się może zdarzyć.

Uświadomiona niewiedza – niepewność relacji jest tak duża, że sformułowanie scenariuszy jest niemożliwe.

Natura niepewności

Natura jest wymiarem pozwalającym odróżnić niepewność wynikającą z niepełnej wiedzy o analizowanym zjawisku od naturalnej zmienności charakteryzującej badane zjawisko.

To rozróżnienie jest ważne, gdyż różne rodzaje niepewności wymagają odmiennego traktowania. Walker i in. (2003) używają określenia *variability* dla niepewności wynikającej ze zmienności czy przypadkowości i *epistemic* dla określenia niepewności związanej z niedostateczną wiedzą.

Inne terminy określające naturę niepewności spotykane w publikacjach (głównie anglojęzycznych) to: w pierwszym przypadku: *type I, type A, aleatory, ontological*; w drugim: *type II, type B* (por. Kowalczyk 2010).

Niepewność wiedzy, może być zmniejszona poprzez dokładniejsze zbadanie zjawiska. Identyfikacja tego typu niepewności w jakimś obszarze modelowania, w połączeniu z oceną jej wpływu na otrzymywane wyniki, daje podstawy do wytyczania kierunków doskonalenia procesu progностycznego. Pozwala np. stwierdzić, czy wysiłek powinien być skoncentrowany na estymacji parametrów czy może na lepszym określeniu wejścia do modelu.

W literaturze ekonomicznej nt. prognozowania najwięcej uwagi poświęca się lokalizacji niepewności. (Np. Clements i Hendry (1995) analizują ten wymiar bardzo szczegółowo, czyniąc z niego podstawę klasyfikacji błędów prognoz).

Problem natury niepewności jest wprawdzie obecny w świadomości ekonomistów od dawna dzięki książce Knighta (1921), ale świat modeli ekonometrycznych (a także modeli matematyki finansowej) zajmuje się tylko jednym z aspektów opisywanych przez Kni-

gha. Knight wskazywał na istnienie dwóch kategorii niepewności, z którymi spotykają się podmioty gospodarcze. Nazywał je odpowiednio ryzykiem i niepewnością, zwracając uwagę na istotne rozróżnienie pomiędzy mierzalnym, możliwym do wyceny ekonomicznej ryzykiem a niepewnością towarzyszącą przewidywaniom. Warunkiem mierzalności wg Knighta jest możliwość określenia obiektywnego prawdopodobieństwa wystąpienia zdarzenia czy rezultatu (metodami statystycznymi lub poprzez wyspecyfikowanie wszystkich jednakowo prawdopodobnych wyników i ich odpowiednie grupowanie). Pisząc o niemierzalnej niepewności Knight posługiwał się prawdopodobieństwem subiektywnym, zaznaczając jednocześnie, że przy jego określaniu może być celowe wykorzystanie wiedzy na temat znanych rozkładów obiektywnych.

Zaproponowany przez Knighta podział niepewności na mierzalną i niemierzalną pokrywa się znaczeniowo z podziałem przeprowadzonym przez Walkera i in. na *variability* i *epistemic uncertainty* (przez *variability* należy rozumieć zarówno zmienność, jak i przypadkowość).

Konstruując wykresy wachlarzowe w oparciu o rozkłady opisujące błędy statystyczne lub zaobserwowaną w przeszłości zmienność, abstrahujemy od niepewności w rozumieniu Knighta, która z punktu widzenia odbiorców prognoz makroekonomicznych (gremiów decyzyjnych czy podmiotów gospodarczych) może być bardzo istotna.

3. Drugi wymiar niepewności w popularnych metodach

Prosty sposób łączenia subiektywnych opinii ekspertów (niepewność typu II) z wiedzą dostępną dzięki modelom (związki przyczynowe) zdecydowało o popularności metody opracowanej w Banku Szwecji (Blix i Sellin 1999, 2000). Pierwsze wykresy wachlarzowe publikowane przez NBP były również tworzone w oparciu o tę metodę.

Metoda Blixa i Sellina wykorzystuje mnożniki modelowe i zależności między wariancjami w modelach liniowych, co powoduje, że jej zastosowanie nie zawsze jest możliwe. Dosyć uniwersalny wydaje się natomiast sposób konwersji opinii ekspertów na rozkłady prawdopodobieństwa.

Do opisu przyszłej inflacji (metoda była też stosowana dla pkb) a także zmiennych, od których ona zależy, używa się rozkładów binormalnych (*binormal*, *two-piece normal*, *TPN*) a więc ich gęstość ma następującą postać:

$$f(x) = \begin{cases} A \exp(-(x - \mu)^2 / 2\sigma_1^2) & \text{dla } x \leq \mu \\ A \exp(-(x - \mu)^2 / 2\sigma_2^2) & \text{dla } x > \mu \end{cases} \quad (1)$$

gdzie $A = \sqrt{2/\pi}(\sigma_1 + \sigma_2)^{-1}$.

Niech $\{Z_j\}_{j=1,\dots,n}$ będzie zbiorem zmiennych determinujących inflację.

Parametry rozkładu inflacji wyznaczone są w oparciu o opinie ekspertów dotyczące przyszłych wartości zmiennych Z_j , $j = 1, \dots, n$. Eksperti podają dla każdej z nich wartość najbardziej prawdopodobną $\mu_j(t)$ i określają dla kolejnych horyzontów:

- 1) prawdopodobieństwo $P_j(t)$ wystąpienia wartości niższych od $\mu_j(t)$;
- 2) stopień wzrostu/spadku niepewności $h_j(t)$ w porównaniu z niepewnością historyczną, co pozwala obliczyć odchylenie standardowe $\sigma_j(t) = h_j \sigma_{hist}(t)$.

W oparciu o te informacje obliczane są parametry $\sigma_{j,1}(t)$ i $\sigma_{j,2}(t)$ rozkładu binormalnego opisującego zmienną Z_j w horyzoncie t :

$$\sigma_{j,1}^2(t) = \sigma_j(t) \left[(1 - 2/\pi) \left(\frac{1 - 2P_j(t)}{P_j(t)} \right)^2 + \left(\frac{1 - P_j(t)}{P_j(t)} \right)^2 \right]^{-1} \quad (2)$$

$$\sigma_{j,2}^2(t) = \sigma_j(t) \left[(1 - 2/\pi) \left(\frac{1 - 2P_j(t)}{1 - P_j(t)} \right)^2 + \left(\frac{P_j(t)}{1 - P_j(t)} \right)^2 \right]^{-1}$$

Znajomość $\sigma_{j,1}(t)$ i $\sigma_{j,2}(t)$ pozwala wyznaczyć parametr asymetrii $\gamma_j(t)$. Jest on definiowany jako różnica między wartością oczekiwaną i dominantą (definicja ta generuje pewne problemy, do których wrócimy w rozdz. 6) i może być wyrażony w następujący sposób:

$$\gamma_j(t) = \sqrt{(2/\pi)}(\sigma_{j,2}(t) - \sigma_{j,1}(t)) \quad (3)$$

Asymetria inflacji $\gamma(t)$ jest aproksymowana kombinacją liniową asymetrii zmiennych

Z_j , $j = 1, \dots, n$:

$$\gamma(t) = \sum_{j=1}^n \beta_j(t) \gamma_j(t) \quad (4)$$

Współczynnikami są elastyczności β_j otrzymane przy pomocy modelu.

Dodatnie wartości $\gamma(t)$ wskazują na większe ryzyko wystąpienia wartości wyższych od wartości najbardziej prawdopodobnej (*upside risk*), ujemne - niższych (*downside risk*).

Znajomość parametru asymetrii $\gamma(t)$ pozwala na znalezienie parametrów $\sigma_1(t)$ i $\sigma_2(t)$ rozkładu opisującego inflację (wcześniej na podstawie odpowiedzi ekspertów i błędów historycznych szacuje się wariancję $\sigma^2(t)$).

Blix i Sellin nazywają otrzymywany w ten sposób rozkład „częściowo subiektywnym” podkreślając w ten sposób to, że punktem wyjścia są rozkłady obiektywne opisujące przeszłe błędy.

Wykorzystywanie rozkładów błędów historycznych w charakterze rozkładów bazowych i modyfikowanie ich w oparciu o informacje spoza próby jest standardowym podejściem stosowanym przy tworzeniu wykresów wachlarzowych, niezależnie od tego czy stosowane są metody analityczno-aproksymacyjne czy też symulacje. Istotną wadą takiego podejścia jest brak identyfikowalności składowych niepewności o odmiennym charakterze. Do problemu powrócimy w rozdziale 5 a w rozdziale 6 zostanie przedstawiona propozycja rozwiązania.

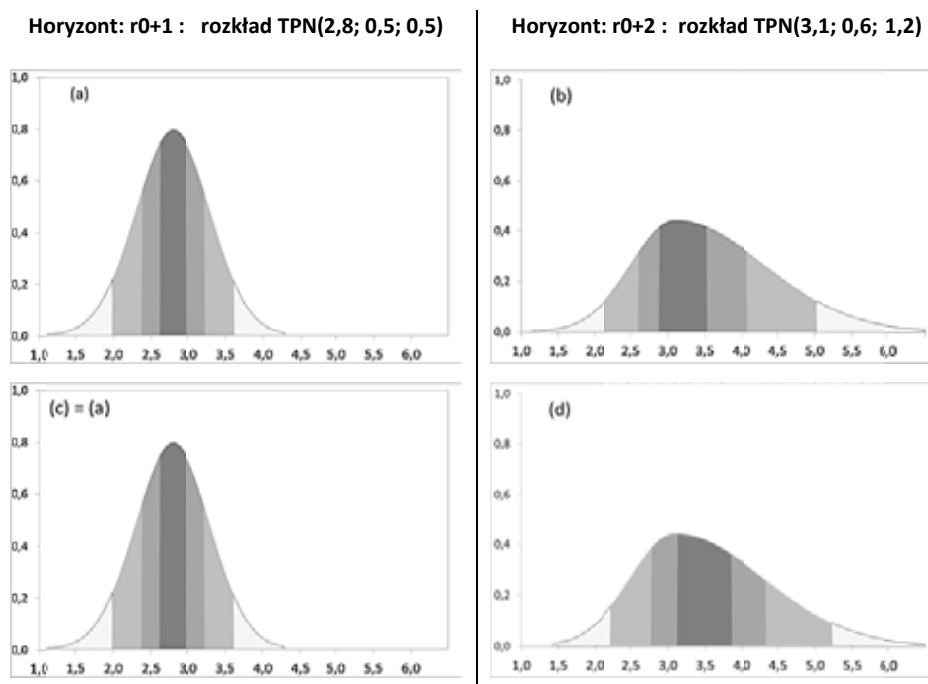
4. Trzeci wymiar niepewności a dotychczasowe metody

Problem skali niepewności nie był dotąd brany pod uwagę przy konstrukcji wykresów wachlarzowych. Przyjmowane jest założenie, że istnieje pewien wyróżniony scenariusz, który jest bardziej prawdopodobny od innych. Znajduje to odbicie w postaci zakładanych rozkładów - są one jednodymalne. Sygnalizowanie możliwości wystąpienia odmiennych scenariuszy sprowadza się zazwyczaj do wprowadzenia dużej asymetrii rozkładu. Może to prowadzić do problemów przy przekazywaniu informacji o niepewności, jeżeli przedziały prawdopodobieństwa są budowane wokół median, a nie wokół dominant.

Zilustrujemy to rozpatrując następującą hipotetyczną sytuację. Załóżmy, że prognozując inflację, w wyniku analizy niepewności otrzymaliśmy dla dwóch kolejnych rocznych horyzontów funkcje gęstości przedstawione na wykresach 2 a-d. Pierwsza (horyzont: r_0+1) jest symetryczna, natomiast drugą (horyzont: r_0+2) cechuje duża asymetria.

Wykresy 2 a-d

Funkcje gęstości prawdopodobieństwa dla dwóch kolejnych horyzontów i różne sposoby konstrukcji przedziałów o zadanym prawdopodobieństwie: a,b - przedziały budowane wokół dominanty (najkrótsze); c,d - przedziały budowane wokół mediany



Znając funkcję gęstości możemy znaleźć dowolny kwantyl, co pozwala skonstruować przedziały o zadanym prawdopodobieństwie. Przedziały te nie są jednoznacznie określone. Np. prawdopodobieństwo każdego z przedziałów wyznaczonych przez najciemniejsze obszary na wykresach 2b i 2d jest równe 0,3 a granice przedziałów są różne.

Przedział na wykresie 2b ma tę właściwość, że jest przedziałem o najmniejszej długości spośród wszystkich przedziałów o prawdopodobieństwie 0,3. Wartości funkcji gęstości prawdopodobieństwa na jego końcach są sobie równe. Tak wybierane przedziały są podstawą konstrukcji wykresów wachlarzowych, które nazwiemy modalnymi. Przykładem takich wykresów są wykresy wachlarzowe publikowane przez Bank Anglii.

Z kolei środkowy przedział na wykresie 2d jest przedziałem wyznaczonym przez kwantyle rzędu 0,35 i 0,65. Jest symetryczny (biorąc pod uwagę prawdopodobieństwo a nie długość) względem mediany, która dzieli go na dwa przedziały o prawdopodobieństwie 0,15. Wykresy wachlarzowe budowane wokół mediany rozkładu będziemy nazywać kwantylowymi. Ich zaletą jest to, że pozwalają oszacować prawdopodobieństwa wartości wyższych i niższych od ścieżki centralnej (w przypadku wykresów modalnych znana jest tylko suma prawdopodobieństwa obszarów jednolicie pokolorowanych). Wykresy kwantylowe były wykorzystywane przez Bank Szwecji (do momentu przejścia na błędy statystyczne *ex post*). Występują też w „Raportach o Inflacji” publikowanych przez NBP (ten typ wykresów został wybrany m.in. pod wpływem uwag Wallisa (1999) do projekcji Banku Anglii a także z obawy przed wielomodalnością, która mogła wystąpić przy agregowaniu prognoz otrzymywanych z różnych modeli).

Jak będą wyglądały wykresy wachlarzowe modalne i kwantylowe, pokazujące przedziały o prawdopodobieństwie 0,3; 0,6 i 0,9 skonstruowane na bazie omawianych funkcji gęstości, ze ścieżką centralną pokazującą scenariusz najbardziej prawdopodobny, przedstawiono na wykresach 3a-3b (odchylenia standardowe dla kwartałów 1-3 w kolejnych latach otrzymano drogą interpolacji zakładając równomierny wzrost niepewności w ciągu roku).

Oba typy wykresów wyraźnie sygnalizują większe prawdopodobieństwo wystąpienia wyższych wartości, ale w przypadku wykresu kwantylowego niewiele brakuje, by ścieżka centralna znalazła się poza pasmem centralnym, co mogłoby rodzić zastrzeżenia dotyczące spójności przekazu.

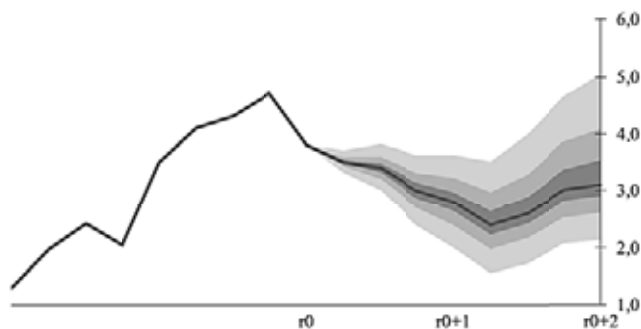
Na wykresach modalnych ścieżka centralna leży zawsze w obszarze najciemniejszym. Można zrezygnować z jej przedstawiania (by odbiorcy prognoz nie przywiązywali zbytniej uwagi do prognoz punktowych) bez obaw, że ucierpi na tym przekaz dotyczący asymetrycznego ryzyka. Inaczej jest w przypadku wykresów kwantylowych. Pominięcie ścieżki najbardziej prawdopodobnej lub zastąpienie jej ścieżką wartości oczekiwanych, spowoduje, że asymetria nie będzie wystarczająco sygnalizowana, jak to pokazano na rysunku 3c.

Wykresy 3 a-c

Silna asymetria a różne typy wykresów wachlarzowych: a-modalny; b, c - kwantylowy. Rola ścieżki centralnej: a, b - ścieżki centralne wyznaczone są przez dominanty rozkładów; c - kwantylowy ze ścieżką określoną przez wartości oczekiwane

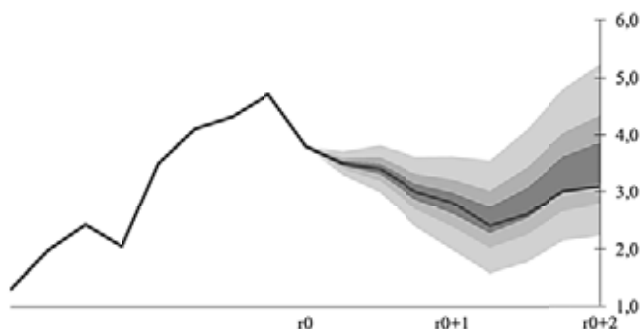
a

modalny/dominanty



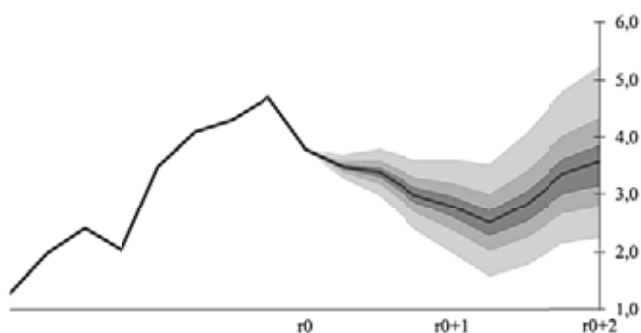
b

kwantylowy/dominanty



c

kwantylowy/wart.oczek.



Choć wykres 3c różni się od 3b tylko typem ścieżki centralnej, silna asymetria nie jest widoczna. Jest jeszcze inny problem, który wydaje się istotny w przypadku, gdy poprzez wykres wachlarzowy chcemy przekazać opinie nt. przyszłych zjawisk - pominięcie informacji o dominancie, czyli hipotezie uznanej za najbardziej prawdopodobną, może być trudne do zaakceptowania.

Przedstawiając różne typy wykresów wachlarzowych pokazaliśmy, że problemy z przekazem informacji o niepewności mogą wystąpić nawet wtedy, gdy poziom niepewności typu II nie przekracza statystycznego (w rozumieniu Walkera i in).

Standardowe, bazujące na rozkładach jednomodalnych, metody tworzenia wykresów wachlarzowych wydają się wystarczające w okresach względnej stabilności - można wtedy poprzestać na przedstawianiu „najlepszej” predykcji (bez względu na to, która z charakterystyk punktowych zostanie przyjęta) i na lokalnej analizie niepewności polegającej na opisywaniu odchyłeń od ścieżki centralnej. Zawodzą, gdy mamy do czynienia z dużą asymetrią lub gdy w ogóle nie można zakładać jednego dominującego scenariusza i konieczne staje się dostarczenie odbiorcom prognoz informacji o potencjalnych konsekwencjach przyjęcia innych założeń lub podjęcia odmiennych decyzji.

5. Problem wymieszania różnych typów niepewności

W sytuacji, gdy wykresy wachlarzowe tworzone są na bazie zmodyfikowanych rozkładów błędów historycznych (jest to typowe dla większości stosowanych metod), powstaje problem z oceną wpływu poszczególnych składowych niepewności na rozkład wynikowy. Odbiorcy prognoz na ogół nie otrzymują informacji, jaka część wariancji (decydującej o szerokości przedziałów na wykresie wachlarzowym) wynika z niedoskonałości metod prognozytycznych, a jaka z niepewności co do możliwych scenariuszy. Ponadto, jeżeli błędy historyczne są duże, przekaz dotyczący asymetrycznych odchyień od ścieżki centralnej może ulec zniekształceniu, gdy za miarę asymetrii przyjmie się różnicę wartości oczekiwanej i dominanty. (Więcej na ten temat w kolejnym rozdziale).

Problem znoszenia asymetrii jest często obserwowany, gdy do wyznaczania rozkładów stosowane są symulacje polegające na wielokrotnym rozwiązywaniu modelu – dla różnych zestawów danych wejściowych. W symulacjach *Monte Carlo* dane są losowane zgodnie z założonymi rozkładami prawdopodobieństwa (poważnym problemem jest właściwe odтворzenie zależności rozkładów). W przypadku nieparametrycznych technik bootstrapowych losuje się bezpośrednio z próby. Liczba przeprowadzonych symulacji musi być wystarczająco duża, by dać podstawy do wyznaczenia parametrów rozkładów a w konsekwencji - uzyskać rozwiązania i przedziały niepewności.

Symulacje pozwalają uwzględnić wiele źródeł niepewności, ale mogą też prowadzić do wymieszania i nierozróżnialności poszczególnych składowych, co sprawia, że otrzymywane rozkłady są trudne do zinterpretowania.

Paradoksalnie, dążenie do otrzymania pełnego obrazu niepewności może prowadzić do zmniejszenia wartości informacyjnej prognoz i ich przydatności w procesie decyzyjnym, jeżeli nie uwzględnimy jej wielowymiarowego charakteru. W wielu publikacjach podkreśla się, że konieczne jest odseparowanie skutków niepewności wiedzy od zmienności (np. Hoffman, Hammonds 1994; Frey, Burmaster 1999; Wu, Tsang 2004). Proponowane są np. symulacje *Monte Carlo drugiego rzędu (second order, two-phase, two-dimensional)* polegające na osobnej propagacji tych różnych typów niepewności.

W przypadku prognoz makroekonomicznych ważne wydaje się odseparowanie skutków niepewności samego modelu (można ją traktować jako mierzalną) od niepewności związanej z obszarem wejścia do modelu (punkt startowy, przyszłe wartości zmiennych egzogenicznych). Do opisanie niepewności modelu mogą być wykorzystane obiektywne rozkłady

zaobserwowanych błędów czy rezyduów. Opisanie niepewności zlokalizowanej w obszarze wejścia będzie na ogół wymagało odwołania się do opinii ekspertów i użycia rozkładów subiektywnych związanych z personalistyczną a nie częstościową interpretacją prawdopodobieństwa. (Szerokie omówienie różnych interpretacji można znaleźć np. w Hájek 2009; w literaturze polskojęzycznej - Szreder 2004).

6. Propozycja innego podejścia do konstrukcji wykresów wachlarzowych

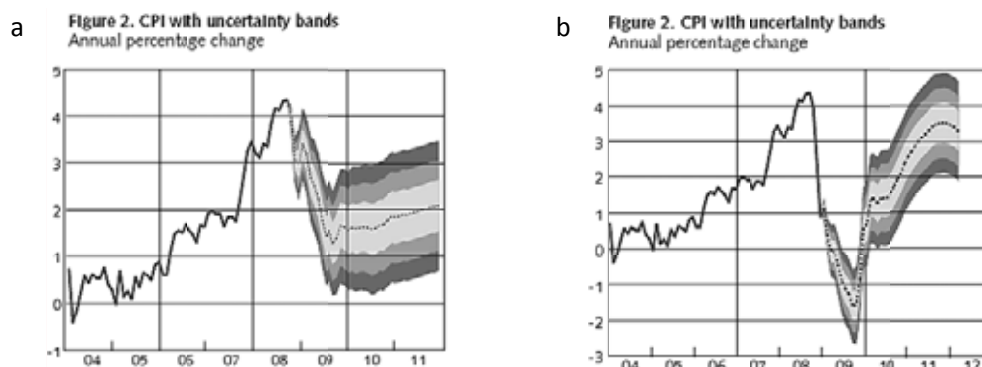
Przejdziemy teraz do przedstawienia propozycji nieco innego sposobu konstrukcji wykresów wachlarzowych. Zaproponujemy metodę, która co prawda zawiera zarówno elementy statystyczne jak i eksperckie, ale pozwala na osobną analizę składowych niepewności o odmiennym charakterze. Jest to możliwe dzięki zastosowaniu łączenia rozkładów, podczas gdy tradycyjne metody polegają na modyfikacji parametrów rozkładów.

Zostaną wyeliminowane problemy ujawniające się w sytuacjach, gdy niepewność zlokalizowana w obszarze wejścia do modelu jest zbyt duża, by można było poprzestać na opisywaniu prawdopodobieństwa odchylenia od ścieżki centralnej i konieczne jest rozważenie alternatywnych scenariuszy.

Z taką dużą niepewnością mieliśmy do czynienia np. na przełomie 2008 i 2009 roku. Wykresy 4 a-b z dwóch kolejnych edycji *Monetary Policy Report* Banku Szwecji z tego okresu są sugestywnym przykładem, jak niewiele wynika z wykresów wachlarzowych pokazujących jeden scenariusz i błąd statystyczny. Zmiana przebiegu ścieżki centralnej powoduje przesunięcie „przywiązanego” do niej wachlarza wraz z nią.

Oczywiście nie znamy odpowiedzi na pytanie: czy scenariusz odpowiadający ścieżce centralnej z wykresu 4b mógł być uznany kilka miesięcy wcześniej za możliwy. Bez studiowania dodatkowych materiałów (np. przedstawianych przez Bank Szwecji aktualizacji) możemy jedynie stwierdzić, że w okresie pomiędzy raportami musiała pojawić się niepewność znacznie większa niż statystyczna, co nie zostało odzwierciedlone w projekcjach.

Wykresy 4 a-b Projekcje inflacji z błędami statystycznymi w okresach dużej niepewności



Źródła : a - „Monetary Policy Report 2008:3”; b - “Monetary Policy Report February 2009”; Sveriges Riksbank

Istotą proponowanego podejścia jest dwuetapowe – hierarchiczne wyznaczanie funkcji gęstości prawdopodobieństwa, tak by uwzględnić zarówno niepewność co do możliwych scenariuszy, jak i to, że każdemu z nich towarzyszy nieunikniony błąd statystyczny. Nie będziemy zakładać istnienia scenariusza dużo bardziej prawdopodobnego od pozostałych ani postaci rozkładu wynikowego.

Konstrukcja wykresu wachlarzowego będzie wymagała określenia, dla każdego z analizowanych horyzontów, dwóch rozkładów prawdopodobieństwa: pierwszego o charakterze subiektywnym, drugiego – obiektywnym.

Rozkład pierwszy, związany z wejściem do modelu, będzie odzwierciedlał opinie ekspertów nt. możliwości realizacji różnych scenariuszy ekonomicznych (istotnych z punktu widzenia procesu decyzyjnego). Przypisane scenariuszom wagi będą implikowały rozkład wartości zmiennych wyjściowych (np. inflacji, pkb). Może to być rozkład dyskretny, gdy liczba analizowanych scenariuszy jest określona (np. przy symulacjach deterministycznych ograniczonych do kilku wariantów) lub ciągły, gdy do wyników analiz czy symulacji scenariuszowych dopasowywana jest funkcja gęstości.

Pojęcie scenariusza nie musi być oczywiście ograniczone do różnych wartości zmiennych egzogenicznych czy punktu startowego – może również dotyczyć przewidywanych odstępstw od wcześniej obserwowanych zależności.

Rozkład drugi będzie opisywał statystyczną niepewność modelu (ogólniej: systemu prognostycznego). Może być tworzony na bazie błędów historycznych oczyszczonych z błędów wynikających z założonych danych wejściowych (by nie duplikować niepewności z tych samych źródeł). W przypadku modeli estymowanych naturalne byłoby też wykorzystanie reszt równań.

Dyskretny rozkład subiektywny

Niech $\{x_i(t)\}_{i=1\dots N}$ będzie zbiorem ścieżek inflacji (prognoz punktowych dla kolejnych okresów) odpowiadających możliwym scenariuszom ekonomicznym. Szanse na realizację i – tego scenariusza oznaczmy przez p_i .

Chcielibyśmy uwzględnić niepewność modelu dla prognozy dotyczącej horyzontu t_k (argument t_k zostanie pominięty dla uproszczenia zapisu). Potraktujemy model jako swego rodzaju miernik służący do wykonywania pomiarów pośrednich, tzn. „zmierzone” dane wejściowe (w naszym przypadku założone wartości zmiennych składających się na okre-

ślony scenariusz) są przetwarzane wg algorytmu określonego równaniami modelu na wynik pomiaru - wielkości wyjściowe (np. inflację, pkb).

Zatem: uwzględnienie błędu „pośredniego pomiaru” inflacji wynikającego z niedoskonałości modelu, wymaga zastąpienia deterministycznej zmiennej x_i , zmienną $X_i = x_i + \Delta X_i$, gdzie ΔX_i jest losowym błędem pomiaru.

Jak wspomnieliśmy, dokładność „miernika” (błąd modelu) może być oceniona metodami statystycznymi, a efekty propagacji błędu dla kolejnych horyzontów np. poprzez symulacje. Skoncentrujemy uwagę na ustalonym horyzoncie prognozy.

Błąd „pomiaru” dotyczy w jednakowym stopniu każdego scenariusza, więc przyjmiemy, że ΔX_i ($i = 1, 2, \dots, N$) są zmiennymi losowymi o identycznym rozkładzie. Oznaczmy go przez $g(x)$. Rozkład obciążonej błędem prognozy X_i odpowiadającej i -temu scenariuszowi będzie wtedy opisany rozkładem:

$$g_i(x) = g(x - x_i) \quad (5)$$

Znając rozkłady p i g możemy uzyskać rozkład będący probabilistycznym opisem następującej sytuacji: w pierwszym kroku z prawdopodobieństwem p_i wybieramy scenariusz; w drugim losujemy wartość błędu i otrzymujemy wartość wskaźnika.

Taki hierarchiczny sposób postępowania prowadzi do rozkładu mieszanego, którego funkcję gęstości oznaczmy przez f :

$$f(x) = \sum_i^S p_i g(x - x_i) \quad (6)$$

Ciągły rozkład subiektywny

Gdy do opisu możliwych scenariuszy użyjemy ciągłego rozkładu $p(s)$, zamiast (6) otrzymamy:

$$f(x) = \int_{-\infty}^{+\infty} p(s) g(x - s) ds \quad (7)$$

Wybór scenariusza następuje zgodnie z subiektywnym rozkładem p , tzn. $p(s)$ jest wagą scenariusza prowadzącego do wartości s . Uwzględniając błąd dla ustalonego s otrzymamy wartość wynikającą z rozkładu $g(x - s)$.

Dla błędu opisanego rozkładem normalnym:

$$f(x) = (1/\sqrt{2\pi\sigma_g^2}) \int_{-\infty}^{+\infty} p(s) \exp(-(x-s)^2 / 2\sigma_g^2) ds \quad (8)$$

Obliczenie $f(x)$ będzie na ogół wymagało całkowania numerycznego. Postać analityczną możemy łatwo otrzymać np. w przypadku, gdy g i p są rozkładami normalnymi. Rozwiązanie dla $p = TPN$ przedstawimy w dalszej części rozdziału.

Oto kilka sytuacji, gdy utworzenie ciągłego rozkładu $p(s)$ wydaje się dosyć proste i uzasadnione:

- Jeżeli w wyniku przeprowadzonych analiz scenariuszowych jesteśmy w stanie określić wartość wskaźnika odpowiadającą scenariuszowi pesymistycznemu, najbardziej prawdopodobnemu i optymistycznemu, niepewność scenariuszową możemy odzwierciedlić przy pomocy rozkładu trójkątnego.
- Jeżeli możemy założyć jednomodalność oraz podać wartość najbardziej prawdopodobną i prawdopodobieństwa wystąpienia wartości od niej niższych lub wyższych, do wykorzystania jest rozkład binormalny (TPN).
- Jeżeli możemy podać wartość wskaźnika odpowiadającą scenariuszowi centralnemu i przedziały o określonym prawdopodobieństwie (np. 50-cio i 90-cio procentowe) to można utworzyć rozkład „kawałkami jednostajny”.

Zauważmy, że gęstość $f(x)$ dana wzorem (7) jest splotem gęstości, czyli gęstością rozkładu sumy dwóch niezależnych zmiennych losowych o rozkładach p i g . Jest to zgodne z naszą wcześniejszą interpretacją gęstość f jako rozkładu opisującego wynik pomiaru zmiennej X o rozkładzie p , gdy błąd pomiaru err ma rozkład g . Wynik pomiaru jest wtedy sumą $X + err$ o rozkładzie opisanym splotem gęstości.

Szczególny przypadek ciągłego rozkładu subiektywnego – TPN

Wobec popularności rozkładu binormalnego w obszarze komunikowania niepewności projekcji inflacji i pkb, rozważymy sytuację, gdy $p = TPN(\mu_p, \sigma_{1p}, \sigma_{2p})$ a $g = N(0, \sigma_g)$.

Rozkład normalny jest szczególnym przypadkiem rozkładu binormalnego. Zatem splot

$p * g$ zapiszemy w postaci:

$$p * g = TPN(\mu_p, \sigma_{1p}, \sigma_{2p}) * TPN(0, \sigma_g, \sigma_g) \quad (9)$$

co pozwoli na skorzystanie z rezultatów opisanych w Garvin i McClean (1997) dotyczących postaci i parametrów rozkładu sumy niezależnych zmiennych losowych o rozkładach binormalnych.

Po pierwsze: możemy przyjąć, że splot rozkładów binormalnych należy do klasy rozkładów binormalnych.

Po drugie: trzeci moment centralny sumy dwóch zmiennych o rozkładzie binormalnym jest równy sumie trzecich momentów centralnych składowych.

W naszym przypadku oznacza to, że

$$\bar{m}_3(p * g) = \bar{m}_3(p) + \bar{m}_3(r) = \bar{m}_3(p) \quad (10)$$

gdzie: przez $\bar{m}_3(\dots)$ oznaczyliśmy trzecie momenty centralne dla odpowiednich rozkładów.

Korzystając ze znanych (często cytowanych za Johnem (1982)) wyrażeń dla momentów centralnych rozkładu binormalnego, $\bar{m}_3(p)$ a co za tym idzie $\bar{m}_3(p * g)$, możemy obliczyć w następujący sposób:

$$\bar{m}_3(p * g) = \sqrt{(2/\pi)}(\sigma_{2p} - \sigma_{1p}) \left[\left(\frac{4}{\pi} - 1 \right) (\sigma_{2p} - \sigma_{1p})^2 + \sigma_{2p} \sigma_{1p} \right] \quad (11)$$

Ponadto, wobec niezależności X i err , otrzymujemy:

$$E(X + err) = E(X) = \mu_p + \sqrt{(2/\pi)}(\sigma_{2p} - \sigma_{1p}) \quad (12)$$

$$Var(X + err) = \sigma_p^2 + \sigma_r^2 \quad (13)$$

Znając trzeci moment centralny i wariancję $Var(X + err)$, którą oznaczmy przez σ_{p*g}^2 , możemy wyznaczyć skośność splotu:

$$\frac{\bar{m}_3(p * g)}{\sigma_{p*g}^3} \quad (14)$$

Posłuży ona do przeprowadzenia dekompozycji σ_{p*g}^2 na „wariancję lewostronną” i „prawostronną”, które oznaczmy odpowiednio przez $\sigma_{p*g,1}^2$ i $\sigma_{p*g,2}^2$.

Dekompozycja zostanie przeprowadzona w sposób analogiczny, jak to czynią Blix i Sellin (1999), jednak z wykorzystaniem rzeczywistej skośności, a nie przybliżenia skośności różnicą wartości oczekiwanej i dominanty, która nie jest standaryzowana i nie uwzględnia wartości wariancji.

Punktem wyjścia będzie relacja zachodząca między wariancją σ^2 i parametrami σ_1 i σ_2 w rozkładzie binormalnym $TPN(\mu, \sigma_1, \sigma_2)$:

$$\sigma^2 = (1 - 2/\pi)(\sigma_2 - \sigma_1)^2 + \sigma_1 \sigma_2 \quad (15)$$

Skośność może być obliczona jako $\frac{\bar{m}_3}{\sigma^3}$, albo przy pomocy współczynnika Pearsona, który

dla rozkładu binormalnego jest równy: $\sqrt{(2/\pi)}(\sigma_2 - \sigma_1)/\sigma$.

Obie miary skośności dla umiarkowanie asymetrycznych rozkładów są w przybliżeniu równe (por. Garvin i McClean 1997). (Rozkład p^*g pomimo silnej asymetrii p , będzie na ogół rozkładem o niewielkiej asymetrii ze względu na symetrię g).

Z przyrównania obu miar asymetrii otrzymujemy następującą zależność:

$$\sqrt{(2/\pi)}(\sigma_2 - \sigma_1) = \frac{\bar{m}_3}{\sigma^2} \quad (16)$$

Wykorzystanie zależności (15) i (16) dla splotu p^*g , który również jest rozkładem binormalnym, pozwala obliczyć niewiadome $\sigma_{p^*g,1}^2$ i $\sigma_{p^*g,2}^2$ poprzez rozwiązanie następującego układu równań:

$$\begin{cases} \sqrt{(2/\pi)}(\sigma_{p^*g,2} - \sigma_{p^*g,1}) = \frac{\bar{m}_3(p^*g)}{\sigma_p^2 + \sigma_g^2} \\ \sigma_p^2 + \sigma_g^2 = (1 - 2/\pi)(\sigma_{p^*g,2} - \sigma_{p^*g,1})^2 + \sigma_{p^*g,1}\sigma_{p^*g,2} \end{cases} \quad (17)$$

przy czym: trzeci moment centralny $\bar{m}_3(p^*g)$ dany jest przez (11); drugie równanie w układzie (17) jest równaniem kwadratowym - wybieramy rozwiązanie dodatnie.

Dominanta splotu p^*g może być obliczona w następujący sposób:

$$\mu_{p^*g} = \bar{\mu}_{p^*g} - \sqrt{(2/\pi)}[\sigma_{p^*g,2}(t) - \sigma_{p^*g,1}(t)] = \bar{\mu}_p - \sqrt{(2/\pi)}[\sigma_{p^*g,2}(t) - \sigma_{p^*g,1}(t)] \quad (18)$$

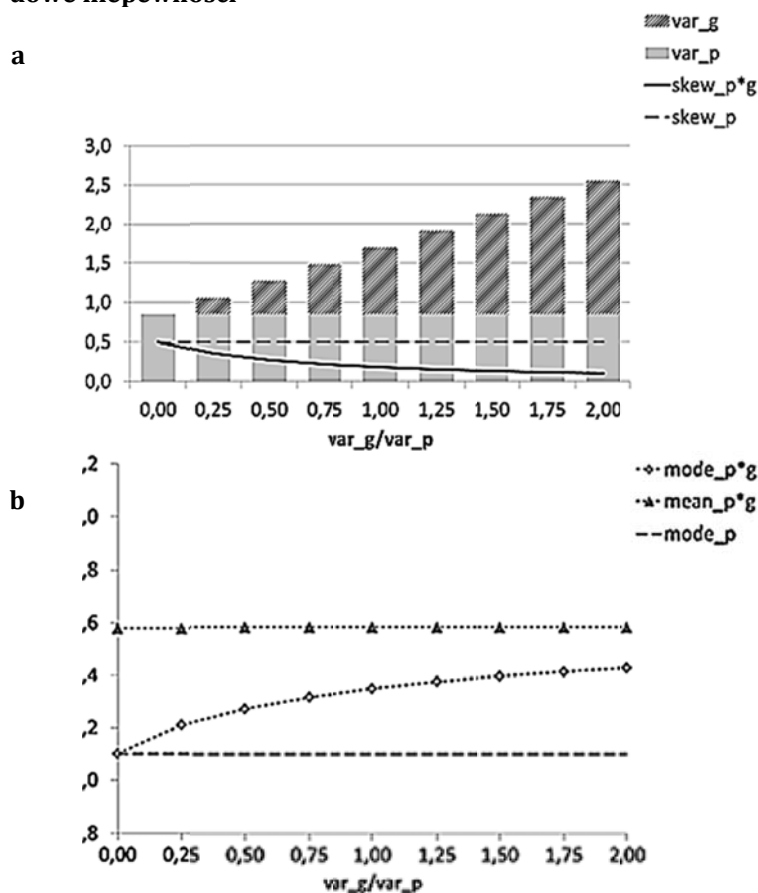
Dominanty μ_{p^*g} i μ_p (w przeciwieństwie do wartości oczekiwanych) są różne.

Warto zwrócić uwagę na to, jak ważna jest informacja o relacji pomiędzy wariancją błędu i wariancją rozkładu opisującego niepewność scenariuszy, której nie dostarczają nam wykresy wachlarzowe konstruowane standardowymi metodami. Duży udział wariancji błędu w wariancji całkowitej powoduje, że rozkład staje się symetryczny, a dominanta przestaje odpowiadać scenariuszowi, który jest najbardziej prawdopodobny.

Wykresy 5a i 5b pokazują, jak zmieniają się skośność i dominanta splotu p^*g w zależności od relacji σ_g^2 do σ_p^2 . Pierwsza obserwacja odpowiada sytuacji, gdy cała niepewność wynika z niepewności danych wejściowych (scenariuszy) a błąd modelu jest zero-owy. W kolejnych udział błędu jest coraz większy.

Wykresy 5a-b

Wpływ rosnącego błędu na: a- skośność, b-dominantę rozkładu opisującego obie składowe niepewności



Założono rozkłady: $p = \text{TPN}(3,1;0,6; ,2)$ (por. wyk. 2) ; $g = N(0,k*\text{var}_p)$, gdzie $k = \text{var}_g/\text{var}_p$;

Objasnienia symboli: var - wariancja; mode - dominanta; mean - wartość oczekiwana;

skew - skośność definiowana w oparciu o momenty;

Rozkład p^*g przedstawiający skumulowaną niepewność nie jest wolny od wad rozkładów otrzymywanych drogą modyfikacji czy symulacji (problem wymieszania niepewności), ale w przeciwieństwie do nich, daje możliwość spójnego przedstawienia informacji, których źródłem jest rozkład subiektywny p (np. poprzez pokazanie jego dominanty i /lub wybranych kwantyli, które odpowiadają rozważanym scenariuszom).

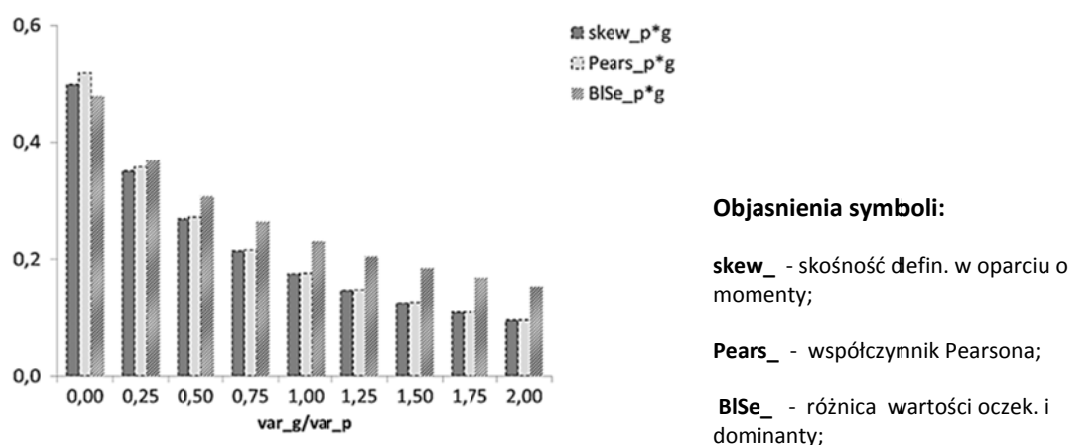
Dzięki wprowadzeniu odrębnych rozkładów charakteryzujących dwa różne aspekty całkowitej niepewności, przekazanie informacji o alternatywnych scenariuszach, jak również ocena wpływu błędu modelu na szerokość przedstawianych przedziałów, nie sprawiają większych trudności. Możliwe jest też analizowanie każdej ze składowych niepewności z uwzględnieniem ich odmiennej natury.

Zwróćmy uwagę na jeszcze jeden problem związany z tradycyjnym podejściem do konstrukcji wykresów wachlarzowych w oparciu o rozkłady binormalne. Otóż asymetria jest zazwyczaj definiowana jako różnica wartości oczekiwanej i dominanty rozkładu (np. Blix, Sellin 1999, 2000) i traktowana jako przybliżenie skośności. Używanie niestandardyzowanej różnicy wartości oczekiwanej i dominanty może prowadzić do zniekształconego obrazu ryzyka asymetrycznych odchyleń, co pokazano na wykresie 6 porównując tę miarę (została oznaczona przez „BISe”) ze współczynnikiem Pearsona i standaryzowanym trzecim moment centralnym obliczonymi dla rozkładów z poprzedniego przykładu.

Konstruując wykres wachlarzowy w oparciu o parametry rozkładu *TPN* wyznaczone na podstawie układu (17) unikamy takich błędów.

Wykres 6

Porównanie różnych miar asymetrii



Przykład wykresu wachlarzowego – dwa alternatywne scenariusze

Założmy, że rozważane są dwa istotnie różniące się scenariusze (pozostałe możemy np. uznać, jako mieszczące się w granicach błędu statystycznego).

Stosując standardowe podejście należałoby skonstruować dwa osobne wykresy wachlarzowe, jak to przedstawiono na wykresach 7a i 7b. Ścieżki alternatywne znalazły się poza obszarami o prawdopodobieństwie 0,9.

Wykresy 7c i 7d są wynikiem zastosowania nowej metody. Rozkłady prawdopodobieństwa zostały wyznaczone w oparciu o funkcje gęstości obliczone zgodnie ze wzorem (6) dla takich samych rozkładów błędów $g_i(x)$.

W pierwszym przypadku przyjęto, że w opinii ekspertów oba scenariusze są jednakowo prawdopodobne. W drugim - szanse na realizację scenariusza sc1 są dwukrotnie wyższe. Prowadzi to do następujących rozkładów subiektywnych:

$$p1: P(X = x_i) = P(S = sc_i) = 1/2 \quad \text{dla } i = 1, 2 \quad (19)$$

$$p2: P(X = x_1) = P(S = sc_1) = 2/3 ; P(X = x_2) = P(S = sc_2) = 1/3 \quad (20)$$

Na każdym z wykresów skonstruowanych wg nowej metody, ścieżki odpowiadające poszczególnym scenariuszom mieszczą się w obszarze o prawdopodobieństwie 0,9.

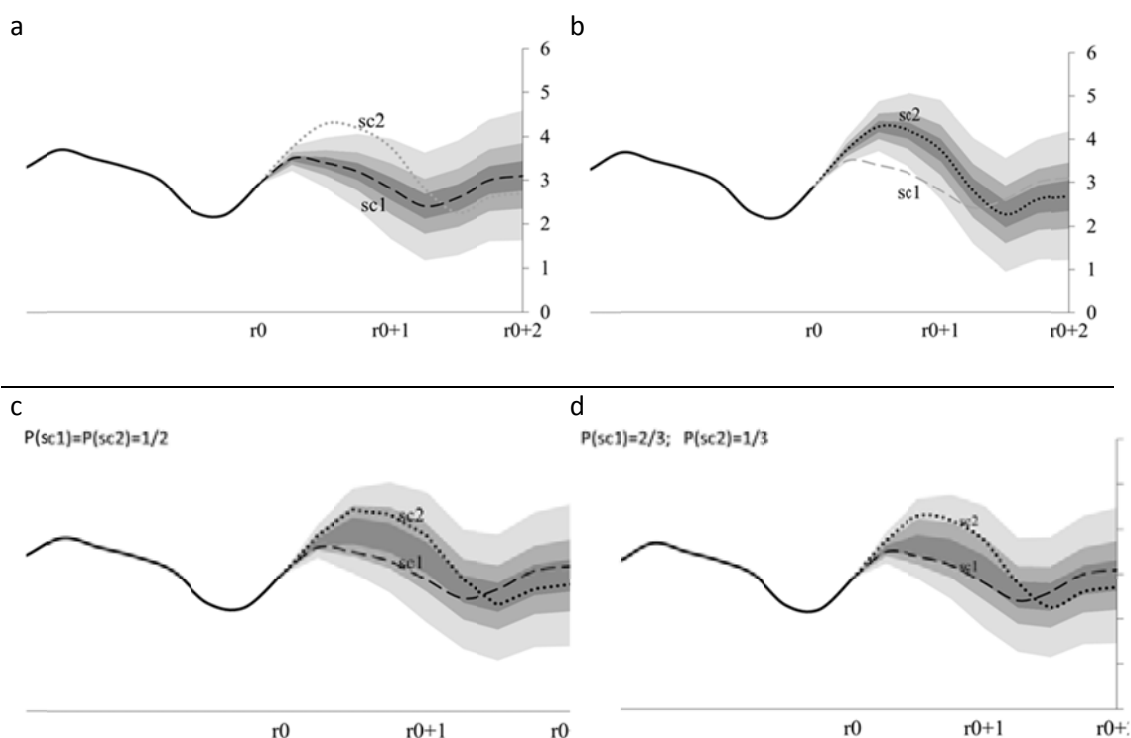
Położenie ścieżek względem pasma centralnego (określonego przez kwantyle) informuje o relatywnym prawdopodobieństwie scenariuszy i pozwala ocenić, czy pasmu centralnemu można przypisać interpretację ekonomiczną czy jedynie statystyczną.

Widać, jaki wpływ na szerokość wachlarza w poszczególnych horyzontach ma niepewność dotycząca scenariuszy.

Wykresy 7 a-d

Porównanie metody standardowej (a,b) i zmodyfikowanej (c,d):

a,b – osobne wykresy dla każdego ze scenariuszy, nie uwzględniona jest niepewność ścieżki alternatywnej; **c** – przypadek gdy oba scenariusze są jednakowo prawdopodobne; **d** – scenariuszowi sc1 przypisano prawdopodobieństwa dwukrotnie wyższe



Na wykresach skonstruowanych wg nowej metody zachowano tradycyjny sposób kolorowania. Wydaje się, że w przypadku rozkładów wielomodalnych wskazane byłoby wprowadzenie zmiany polegającej na uzależnieniu intensywności koloru od wartości funkcji gęstości. Oczywiście należałoby zachować linie odpowiadające parom kwantyli, wyznaczającym przedziały prawdopodobieństwa.

7. Zakończenie

Jeżeli komunikowanie niepewności ma być czymś więcej niż przekazem, że jest ona faktem, z którym należy się pogodzić, to konieczne jest uwzględnienie jej wielowymiarowego charakteru i dostosowanie sposobu prezentacji do potrzeb odbiorców.

Za uniwersalny, istotny z punktu widzenia zarówno twórców, jak i odbiorców prognoz, cel badania i przedstawiania niepewności prognoz można uznać przekazywanie informacji o ograniczeniach związanych ze stosowaniem określonych modeli i technik prognozowania. Zespołom prognostycznym pozwala to wytyczyć kierunki doskonalenia swojej pracy. Odbiorcom pozwala ocenić, w jakim stopniu można ufać przedstawianym prognozom. Wykresy wachlarzowe konstruowane w oparciu o błędy *ex post* spełniają właśnie takie zadanie i choć stają ostatnio dość popularne, to w kontekście decyzyjnym ich przydatność nie jest zbyt duża.

Wiedza o tym, w jakim stopniu można polegać na wykorzystywanych modelach, jest oczywiście niezwykle istotna i powinna wpływać na decyzje o wyborze modelu. Jednak właściwa reakcja na niepewność wymaga, by wykresy wachlarzowe odnosiły się także do możliwych stanów gospodarki i pozwalały zilustrować konsekwencje realizacji różnych scenariuszy czy decyzji. Taki cel przyświecał przedstawionej w niniejszym opracowaniu propozycji modyfikacji dotychczasowych metod, które eksponują jeden scenariusz.

Bibliografia

- Blix, M., Sellin P. (1999), Inflation Forecasts with Uncertainty Intervals, *Sveriges Bank Quarterly Review*, 2.
- Blix, M., Sellin P. (2000), A Bivariate Distribution for Inflation and Output Forecasts, *Sveriges Riksbank Working Paper*
- Britton, E., P. Fisher, Whitley J. (1998), The Inflation Report projections: Understanding the Fan Chart, *Bank of England Quarterly Bulletin* 38
- Charemza W., Jelonek P., Makarova S. (2009), An application of the Choleskized multivariate distributions for constructing inflation 'fan charts', *NBP Workshop 'Experiences and Challenges of Forecasting at Central Banks'*, Warsaw .
- Chatfield Ch. (1995), Model Uncertainty, Data Mining and Statistical Inference, *Journal of the Royal Statistical Society (Series A)*, 158, 419-466.
- Chatfield Ch. (1993), Calculating interval forecasts (with discussion), *J. Bus. Econ. Statist.*, 11, 121-144
- Clements M., Hendry D. (1995), Macro-Economic Forecasting and Modelling, *The Economic Journal* 105, 1001-1013.
- Feldstein M. (1997), The Error of Forecast in Econometric Models when the Forecast-Period Exogenous Variables are Stochastic, *Econometrica* 39, 55-60.
- Fuller W., Hasza D. (1981), Properties of Predictors for Autoregressive Time Series, *Journal of the American Statistical Association*, 76, 155-161.
- Frey H., Burmaster D. (1999), Methods of characterizing variability and uncertainty: comparison of bootstrap simulations and likelihood-based approaches, *Risk Analysis* 19:109-130.
- Garvin J., McClean S. (1997), Convolution and Sampling Theory of the Binormal Distribution as a Prerequisite to Its Application in Statistical Process Control, *Journal of the Royal Statistical Society (Series D)* 46, 33-47.
- Hájek A., Interpretations of Probability, *The Stanford Encyclopedia of Philosophy (Spring 2009 Edition)*, Edward N. Zalta (ed.),
URL=<<http://plato.stanford.edu/archives/spr2009/entries/probability-interpret/>>.

- Hoeting J., Madigan D., Raftery A., Volinsky C. (1999), Bayesian Model Averaging: A Tutorial, *Statistical Science* 14, 382-401.
- Hoffman F., J Hammonds J. (1994), Propagation of uncertainty in risk assessments: the need to distinguish between uncertainty due to lack of knowledge and uncertainty due to variability 14, 707-712.
- John S. (1982), The two-piece normal family of distributions and its fitting, *Communications in Statistical Theory and Methods*, 11, 879-885
- Knight F.H. (1921), *Uncertainty and Profit*, Houghton-Mifflin, Boston
- Kowalczyk H. (2010), O eksperckich ocenach niepewności w ankietach makroekonomicznych, *Bank i Kredyt*, 41, 101-122.
- de Luna X. (2000), Prediction Intervals Based on Autoregression Forecasts, *Journal of the Royal Statistical Society. Series D (The Statistician)*, 49, 87-93.
- Monetary Policy Report 2008:3"*; Sveriges Riksbank.
- Monetary Policy Report February 2009*; Sveriges Riksbank.
- Szreder M. (2004), Od klasycznej do częstościowej i personalistycznej interpretacji prawdopodobieństwa, *Wiadomości Statystyczne*, 8, 1-10.
- Walker W., Harremoës P., Rotmans J., Van der Sluijs J., Van Asselt M., P. Janssen P., Kraymer von Krauss M. (2003), Defining Uncertainty. A Conceptual Basis for Uncertainty Management in Model-Based Decision Support, *Integrated Assessment* 4, 5-17.
- Wallis K. (1999), Asymmetric density forecast of inflation and the Bank of England's fan chart, *National Institute Economic Review*.
- Wu F.-Ch., Tsang Y. (2004), Second-order Monte Carlo uncertainty/variability analysis using correlated model parameters: Application to salmonid embryo survival risk assessment, *Ecological Modelling*, 177(3-4):393-414.