

Larsen, Vegard Høghaug; Thorsrud, Leif Anders

Working Paper

Business cycle narratives

Working Paper, No. 3/2018

Provided in Cooperation with:

Norges Bank, Oslo

Suggested Citation: Larsen, Vegard Høghaug; Thorsrud, Leif Anders (2018) : Business cycle narratives, Working Paper, No. 3/2018, ISBN 978-82-8379-024-5, Norges Bank, Oslo, <https://hdl.handle.net/11250/2500929>

This Version is available at:

<https://hdl.handle.net/10419/210138>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.



<https://creativecommons.org/licenses/by-nc-nd/4.0/deed.no>

WORKING PAPER

Business cycle narratives

NORGES BANK
RESEARCH

3 | 2018

VEGARD H. LARSEN
AND
LEIF ANDERS THORSRUD



NORGES BANK

Working papers fra Norges Bank, fra 1992/1 til 2009/2 kan bestilles over e-post:
FacilityServices@norges-bank.no

Fra 1999 og senere er publikasjonene tilgjengelige på www.norges-bank.no

Working papers inneholder forskningsarbeider og utredninger som vanligvis ikke har fått sin endelige form. Hensikten er blant annet at forfatteren kan motta kommentarer fra kolleger og andre interesserte. Synspunkter og konklusjoner i arbeidene står for forfatterens regning.

Working papers from Norges Bank, from 1992/1 to 2009/2 can be ordered by e-mail
FacilityServices@norges-bank.no

Working papers from 1999 onwards are available on www.norges-bank.no

Norges Bank's working papers present research projects and reports (not usually in their final form) and are intended inter alia to enable the author to benefit from the comments of colleagues and other interested parties. Views and conclusions expressed in working papers are the responsibility of the authors alone.

ISSN 1502-819-0 (online)
ISBN 978-82-8379-024-5 (online)

Business cycle narratives*

Vegard H. Larsen[†]

Leif Anders Thorsrud[‡]

This version
February 23, 2018

Abstract

This article quantifies the epidemiology of media narratives relevant to business cycles in the US, Japan, and Europe (euro area). We do so by first constructing daily business cycle indexes computed on the basis of the news topics the media writes about. At a broad level, the most influential news narratives are shown to be associated with general macroeconomic developments, finance, and (geo-)politics. However, a large set of narratives contributes to our index estimates across time, especially in times of expansion. In times of trouble, narratives associated with economic fluctuations become more sparse. Likewise, we show that narratives do go viral, but mostly so when growth is low. While narratives interact in complicated ways, we document that some are clearly associated with economic fundamentals. Other narratives, on the other hand, show no such relationship, and are likely better explained by classical work capturing the market's animal spirits.

JEL-codes: C55, E32, E71, N10

Keywords: Business cycles, Narratives, Dynamic Factor Model (DFM), Latent Dirichlet Allocation (LDA)

*This Working Paper should not be reported as representing the views of Norges Bank. The views expressed are those of the authors and do not necessarily reflect those of Norges Bank. We thank Knut A. Aastveit, Andre K. Anundsen, Drago Bergholt, Hilde C. Bjørnland, Gunnar Bårdsen, Francesco Furlanetto, and colleagues at BI and Norges Bank for valuable comments. We are grateful to the *Dow Jones Newswires Archive* for sharing their data with us for this research project. This work is part of the research activities at the Centre for Applied Macro and Petroleum economics (CAMP) at the BI Norwegian Business School.

[†]Centre for Applied Macro and Petroleum Economics, BI Norwegian Business School, and Norges Bank. Email: vegard-hoghaug.larsen@bnorges-bank.no

[‡]Corresponding author. Centre for Applied Macro and Petroleum Economics, BI Norwegian Business School, and Norges Bank. Email: leif-anders.thorsrud@norges-bank.no. Norges Bank, P.O. Box 1179, Sentrum 0107, Oslo, Norway.

1 Introduction

In his presidential address before the American Economic Association’s 2017 meeting, Professor Robert J. Shiller writes:

“The human brain has always been highly tuned toward narratives, whether factual or not, to justify ongoing actions,... Narratives “go viral” and spread far, even worldwide, with economic impact...Though these narratives are deeply human phenomena that are difficult to study in a scientific manner, quantitative analysis may help us gain a better understanding of these epidemics in the future.” (Shiller (2017))

This article quantifies the epidemiology of narratives relevant to economic fluctuations, and business cycles in particular, by asking: To what extent are narratives informative for describing business cycle variation, do they go viral, how do they interact with each other, and are they associated with economic fundamentals or better understood as capturing the market’s animal spirits?

To answer these questions, we restrict our attention to narratives told and spread through the mass media, and construct quantitative measures of narratives based on the news topics the media writes about. Shiller (2017) defines the term narrative to mean a simple story or easily expressed explanation of events that many people want to bring up on news. In Section 2 we discuss why the topic modeling approach provides a good quantitative approximation for narratives, while we in Section 3 describe how we technically construct the news topics and transform them into data useful for a time series analysis. We then proceed in four successive steps.

First, in Section 4, we present a daily coincident index model, built to capture aggregate business cycle dynamics, for three major economies; the US, Japan, and Europe (euro area). Unlike conventional models of this type (Stock and Watson (1988), Mariano and Murasawa (2003), Aruoba et al. (2009), and Marcellino et al. (2016)), however, the model allows for time-varying parameters through a threshold mechanism, and, most importantly, uses the daily narratives as input variables. In turn, this innovation allows us to decompose the changes in the latent daily business cycle indexes into time-varying news topic contributions reflecting the continuously evolving narrative about economic conditions, as described by the media. The resulting indexes and decompositions are reported in Sections 4.1 and 4.2.

Building on these results, in Section 4.3, we explore the extent to which narratives relevant for business cycles go viral and affect economic fluctuations and co-movement across borders. In the process we derive novel virality indexes, which provide quantitative and qualitative information about which narratives go viral, when, and for how long.

In Section 4.4 we investigate how narratives independently spread between economic regions. We do so by using the individual news topic time series, their estimated importance for describing business cycle fluctuations, and so called “Graphical Granger causality” modeling (Lozano et al. (2009), Shojaie and Michailidis (2010)). This framework allows us to handle the high dimensionality of the problem, but also draw on graph theory to construct measures of node importance and centrality. More than providing a sophisticated analysis of the causal mechanisms underlying information diffusion, our analysis provides the first attempt of quantifying news spillovers relevant for economic fluctuations for the world’s largest economies.

Finally, in Section 4.5, we show that the complex network of news spillovers can be partitioned into (more or less) exogenous components, and thereby used to cast light on whether narratives are associated with economic fundamentals (Beaudry and Portier (2006), Barsky and Sims (2012), Blanchard et al. (2013)), or noise and sentiment (Shiller (2000), Angeletos and La’O (2013)).

Our analysis is explorative rather than grounded in one (single) formal model. We loosely take a rational inattention view (Sims (2003)), where news broadcasted through the media is important because it can reach a broad population of economic agents and alleviate informational frictions, but also potentially have an independent role in explaining economic fluctuations (Dougal et al. (2012), Peress (2014), Larsen and Thorsrud (2017), Shiller (2017)). We operationalize this view by working with a simple underlying hypothesis: To the extent that the media provides a relevant description of the economy, the more intensive a given topic is represented in the media at a given point in time, the more likely it is that this topic represents something of importance for the economy’s current and future needs and developments. For example, we hypothesize that when the media writes extensively about, e.g., regulatory developments, this reflects that something is happening in this area that potentially has economy-wide effects.

Key to our approach is that we use text as data (Gentzkow et al. (2017)), and our focus on news topics. From the *Dow Jones Newswires Archive* (DJ) we have access to over 40GB of news stories dating back to the early 1990s, covering all areas of economics, a range of countries and regions, and the Dow Jones flagship publication *The Wall Street Journal*.¹ While the Dow Jones news service is far from the monopolistic supplier of economic news, it is among the three biggest suppliers in this global market. Thus, while we can not rightfully argue that we capture all economic news relevant for economic agents in all three countries, we believe the dataset is fairly representative.

The extraction of topics is done using advances in the Natural Language Processing

¹The term “Big Data” is used for textual data of this type because they are, before processing, highly unstructured and contain large amounts of words and articles (Nyman-Andersen (2016)).

literature, while the tone of the news is identified using simple dictionary based techniques (Tetlock (2007)). In general, topic models are statistical algorithms that categorize the corpus, i.e., the whole collection of words and articles, into topics that best reflect the corpus’s word dependencies. In this paper, an unsupervised topic model belonging to the Latent Dirichlet Allocation (LDA) class (Blei et al. (2003)) is used to estimate 80 topics for each country. Each individual topic can be viewed as a word cloud, where the font size used for each word represents how likely it is to belong to this specific topic. We subsequently transform these word clouds into tone adjusted frequency measures, reflecting by how much, and by which tone, each topics is written about on each day in the sample. A vast information set consisting of words and articles can thereby be summarized in a much smaller set of topics facilitating usage in a macroeconomic context. Although topic models hardly have been applied in economic (see, e.g., Hansen et al. (2018) for an exception), their use as a natural language processing tool in other disciplines has been widespread. The LDA’s popularity stems from its success in classifying text and articles into topics in much the same manner as humans would do (Chang et al. (2009)).

We reach five main conclusions. First, in all three countries/regions, the resulting coincident indexes are shown to track the phases of the business cycles with high precision, but performs especially well in the US. For policymakers and forecasters who need to assess the state of the economy in real time to devise appropriate policy responses, the news-based coincident indexes offer a valuable alternative. High-frequency economic statistics covering the broader economy are scarce. Daily news coverage is available in large quantities.

Second, we provide new evidence on the narratives relevant to economic fluctuations. At a broad level, particularly influential news topics include news about macroeconomic developments, the financial market, and (geo-)politics in all three countries. Across time however, there is considerable variation in how narratives contribute to, or describe, economic fluctuations. For example, late in 2007 and through 2008, news about regulatory developments is among the most influential news topics in the US, while earthquake-related narratives became particularly relevant in Japan in 2011. A common pattern across all countries, however, is that in periods associated with recessions, the number of narratives contributing to our index estimates become more sparse than during expansions. Thus, in relation to narratives, expansions are broad based, while recessions are not.

Third, we find that narratives do go viral, as argued by Shiller (2017), but mostly so in times of trouble. In total we identify 13 epidemic episodes between the mid 1995 and 2016, with an average duration of 4-5 months. The narratives contributing the most to these episodes tend to be associated with US-based labor market conditions and (partly)

monetary policy. Interestingly, however, we find little evidence suggesting that epidemics lead to more synchronized international business cycles.

Fourth, the graph describing the network of cross-country news spillovers is dense, but complex. Still, narratives identified with the US dominate, and have predictive power for news in Japan and Europe to a much larger extent than vice versa. The most central nodes in the graphical Granger causality graph are very much the same as those that contribute the most in explaining the fluctuations in the daily coincident indexes, i.e., news about macro economic developments and (geo-)politics, while the least central narratives are found to include news about technology, finance and commodities.

Finally, when partitioning the news topics into more or less exogenous variables using the centrality score computed from the graphical Granger causality graph, we find clear evidence that the most “exogenous” (least connected) narratives are associated with economic fundamentals (total factor productivity (TFP)). Unexpected fluctuations in these narratives lead to persistent, and significant, increases in TFP. In contrast, narratives with a high centrality score show no such relationship. Thus, some narratives confirm to the news-driven business cycle view. Other narratives, on the other hand, are likely better explained by classical work capturing the market’s animal spirits.

This article contributes to a broader literature that seeks to understand the role of narratives in economics (Shiller (2017)). To this end we establish a number of new “stylized facts” about the relationship between business cycles and narratives, epidemics, and cross-country spillovers for the three major economies the US, Japan, and Europe.² As such, we also relate more loosely to a large literature investigating international business cycle synchronization (Kose et al. (2003), Stock and Watson (2005), Mumtaz et al. (2011) Kose et al. (2012)). In contrast to earlier studies in this literature, however, we are the first to focus on narratives. Likewise, by investigating the relationship between narratives and economic fundamentals we speak to a huge and long-lasting literature where changes in economic agents’ expectations, due to either news (new information) or animal spirits (noise/sentiment), are the main underlying driver of business cycle fluctuations (Pigou (1927), Keynes (1936), Beaudry and Portier (2006), Barsky and Sims (2012), Blanchard et al. (2013), Shiller (2000), Angeletos and La’O (2013)).

This paper is also directly related to a large literature, starting with Burns and Mitchell (1946), that seeks to measure business cycles and construct coincident indexes. Regarding the latter, Stock and Watson (1988), Mariano and Murasawa (2003), Aruoba et al. (2009), and Marcellino et al. (2016) provide prominent contributions, and Balke et al. (2017) and

²A closely related field, particularly in finance, investigates the (causal) role of the media itself (Dougal et al. (2012), Peress (2014)). Interestingly, in the first “Handbook of Media Economics” (Simon P. Anderson and Strömberg (2015)) there is a separate chapter about “The Role of Media in Finance” (Tetlock (2015)), but no equivalent chapter about “The Role of Media in Macroeconomics”.

Shapiro et al. (2017) are examples of newer work using text as data. Neither of these studies do, however, provide a narrative account of business cycle fluctuations.

The approach taken here speaks to a growing number of studies in economics using text as data (Bholat et al. (2015), Gentzkow et al. (2017)). On this point, commonly used methods in economics involve some kind of subjectively chosen keyword search and auditing (Baker et al. (2016)), or narrative methods for shock identification Ramey (2016)). For uncovering the narratives relevant for economic fluctuations, the topic modeling approach offers a conceptual advantage over other often applied textual data techniques because it provides interpretable output in a highly automated fashion.³

Lastly, on the methodological side, we draw on recent advances presented in Larsen and Thorsrud (2018) for constructing time series measures of text, and the model proposed in Thorsrud (2016b,a) for constructing the coincident indexes. Both of these studies explore the relationship between news and economic fluctuations in Norway. Here we extend this line of research to three of the biggest economies in the world, and take a narrative perspective. Naturally, we provide a number of news results, and propose new tools for measuring the extent to which narratives go viral, cross-country spillovers, and whether or not narratives are associated with economic fundamentals or animal spirits.

2 On narratives

Humans are inherently storytellers, and the academic literature on narratives is vast. Most work, however, is not found in economic journals, but rather in fields related to linguistics, psychology, anthropology, and history. Here, as alluded to already, we follow Shiller (2017) and define the term narrative to mean a simple story or easily expressed explanation of events that many people want to bring up on news. We then construct measurable approximates to this definition based on the news topics the media writes about, and subsequently link those to economic fluctuations. Accordingly, we will be using the terms narrative and news (topic) interchangeably. More formally, the narrative of a story will consist of one or more news topics. To elaborate on why this approximation is reasonable, what it allows us to measure, and why it might fall short, we take inspiration from the well known cognitive psychologist Jerome Bruner, and in particular Bruner (1991).

First, our interest is not so much in how narratives as text are constructed, but rather how they operate as instruments of mind in the construction, or reflection, of reality.

³For studies that seek to uncover the economic relationships between more concretely defined events or concepts, like, e.g., political uncertainty or monetary policy shocks, a keyword/event search approach might be better suited. For capturing narratives relevant for aggregate business cycles, a keyword/event based approach is not suited unless the researcher knows apriori what to search for.

Obviously, our focus is centered on a narrowly defined aspect of reality, i.e., economic fluctuations, and our sources for constructing measurable narrative approximates are limited to textual news broadcasted through the media. Still, as noted by [Shiller \(2000\)](#); "Significant market events generally only occur if there is similar thinking among large groups of people, and the news media are essential vehicles for the spread of ideas".

We look upon narratives as time dependent, and accounts of events occurring over time. At the same time, "...the particulars of narratives are tokens of broader types" ([Bruner \(1991\)](#)). The modeling approaches adapted in this study reflect these views. As described in greater detail in Section 3, a news story is a weighted sum of different word distributions, i.e., topics. The particular topic composition of a given story, at a given point in time, might very well be unique, but the topic distributions that the narrative constitute are potentially shared by many other narratives. Likewise, to capture the time dependent nature of narratives, we allow the mapping between narratives and economic fluctuations to be time-varying (see Section 4).

However, we do not require the stories in news to be true. Rather, the narrative "truth" is "judged by its verisimilitude rather than its verifiability" ([Bruner \(1991\)](#)). In our setting this means that objective reporting (if that exists) and speculative news stories about market developments, or even news stories about events not happening (if such reporting exists), are all treated equal.

Finally, we take the view that there is only a loose link between the intentional states of a narrative, and the subsequent actions it might induce. Relatedly, the meaning of a story is not simply the sum of its partial expressions, and the interpretation of it will likely depend on the readers background knowledge and context. Admittedly, while neither of these effects are well captured by our approach, it is difficult to envision how quantitative analysis of aggregate economic fluctuation and narratives can fully encapsulate such effects.

3 Data

The main raw data used in this analysis consist of a long sample of daily news extracted from the *Dow Jones Newswires Archive* (DJ). In total we utilize an extraction of over 40GB of raw textual data in XML format from this historical database, which covers a large range of their news services, including content from *The Wall Street Journal*. All text is business-focused, written in English, and covers the US, the Asian, as well as the European market.

The data span the period 1990 to 2016, and includes almost 11 million news articles. Each article listed in the database comes with a number of meta data such as publication

time and region. To classify news as either US, Japan, or Europe specific, we rely on the tags provided by DJ, and partition the dataset accordingly. After removing duplicates and articles that only include updates of earlier published news, the resulting regional data sets include 4754040, 682424, and 1969222 articles for the US, Japan, and Europe, respectively. For all three areas the partitioned data sets end in 2016. For the US we have news observations starting in 1990, while for Japan and Europe the start dates are 1994 and 1995, respectively.

Arguably, what we categorize as country-specific news relies on the DJ definitions, and does not end up as three completely non-overlapping datasets (see Table 16 in Appendix C). As news likely does not stop at the border, we do not find this especially problematic. Another potential limitation is that we have to rely on the DJ region classification tag, and do not use economic news published in region-specific media. As *The Wall Street Journal* is the largest newspaper in the United States in terms of circulation, but likely not in Japan and Europe in general, our raw data might be more representative for the US, than for the two other areas.⁴

To make the textual data applicable for time series analysis, we proceed in three steps illustrated in Figure 1. Technically, these are the same data processing steps as proposed in Larsen and Thorsrud (2018). We provide a summary of the computations below. In the interest of preserving space, technical details are relegated to Appendix C.1 to C.3.

3.1 Cleaning

The share size of the three datasets makes statistical computations challenging. However, as is customary in the Natural Language Processing (NLP) literature, some steps are taken to clean and reduce the raw dataset before estimation (Gentzkow et al. (2017)). First, a stop-word list is employed. This is a list of common words not expected to have any information relating to the subject of an article. Examples of such words are *the*, *is*, *are*, and *this*. In total, the stop-word list together with the list of common surnames and given names removed roughly 1800 unique tokens from the corpus. Next, an algorithm known as stemming is run. The objective of this algorithm is to reduce all words to their respective word stems. A word stem is the part of a word that is common to all of its inflections. An example is the word *effective* whose stem is *effect*. Finally, a measure called *tf-idf*, which stands for term frequency - inverse document frequency, is calculated. This measures how important all the words in the complete corpus are in explaining single articles. The more often a word occurs in an article, the higher the *tf-idf* score of that

⁴Obviously, for us, language barriers are a non-trivial friction in terms utilizing truly country-specific media. Likewise, obtaining textual data of the size and coverage as here is costly. We are grateful to the *Dow Jones Newswires Archive* for sharing their data with us for this research project.

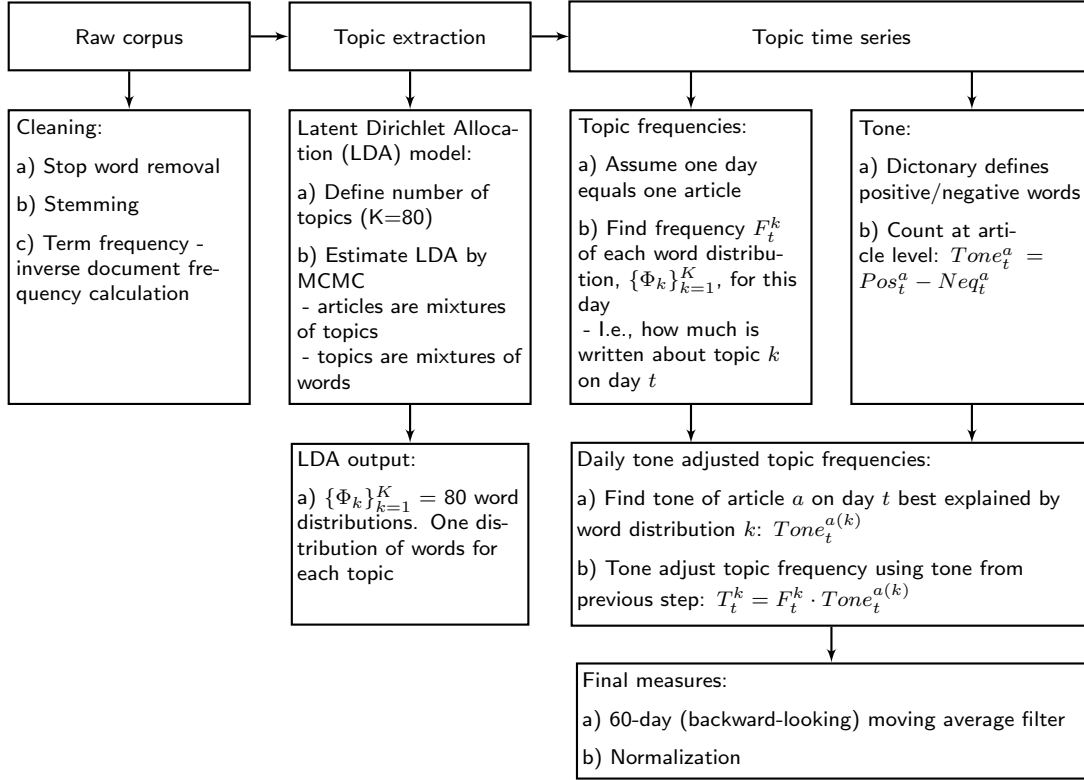


Figure 1. Data preparation flow.

word. On the other hand, if the word is common to all articles, meaning the word has a high frequency in the whole corpus, the lower that word’s *tf-idf* score will be. Around 150 000 of the stems with the highest *tf-idf* score are kept, and used as the final corpus.

3.2 Topic extraction

The “cleaned”, but still unstructured, datasets are decomposed into news topics using a Latent Dirichlet Allocation (LDA) model (Blei et al. (2003)). The LDA model is one of the most popular clustering algorithms in the NPL literature because of its simplicity, and because it has proven to classify text in much the same manner as humans would do (Chang et al. (2009)).

The LDA is an unsupervised topic model that clusters words into topics, which are distributions over words, while at the same time classifying articles as mixtures of topics. A unsupervised learning algorithm is an algorithm that can discover an underlying structure in the data without being given any labeled samples to learn from. The term “latent” is used because the words, which are the observed data, are intended to communicate a latent structure, namely the subject matter (topics) of the article. The term “Dirichlet” is used because the topic mixture is drawn from a conjugate Dirichlet prior.⁵

⁵As such, the LDA shares many features with latent (Gaussian) factor models used in conventional econo-

Different algorithms exist for solving the LDA model. We follow [Griffiths and Steyvers \(2004\)](#), and estimate the model using Gibbs simulations. Technical details and a short description of estimation and prior specifications are described in [Appendix C.1](#). Here we note that we extract $K = 80$ topics from each of the three cleaned datasets. We subjectively chose $K = 80$ for two reasons. First, this was the choice showing the best statistical results in [Larsen and Thorsrud \(2018\)](#) and [Thorsrud \(2016b,a\)](#). Second, we have experimented with estimating both fewer and more topics. It is our experience that with K substantially higher than 80, each topic starts to become highly event specific, i.e., there are signs of over-fitting. Conversely, extracting substantially fewer than 80 topics results in too general topics. Thus, in sum, our choice of $K = 80$ is based on a compromise between fitting the corpus well, getting interpretable topics, as well as earlier experience.

The LDA produces two outputs; one distribution of topics for each article in the corpus, and one distribution of words for each of the topics. Our primary interest is in the latter distributions, which are illustrated using word clouds in [Figure 2](#). Now the LDA estimation procedure does not give the topics any name or label. To do so, labels are subjectively given to each topic based on the most important words associated with each topic. For example, as seen from [Figure 2](#), the most important words associated with the US topic number $T0$ are *monetary*, *inflation*, and *bernanke*. Thus, we label this topic *Monetary Policy*. While it is, in most cases, conceptually simple to classify the topics, the exact labeling plays no material role in the experiment, it just serves as a convenient way of referring to the different topics (instead of using, e.g., long lists of words). A full list of the different topics, their most important words, and our subjective labeling is given in [Tables 9 to 11](#) in [Appendix A](#).⁶

3.3 Topic time series

Given knowledge of the topics (and their distributions), the topic decompositions are translated into tone adjusted time series. To do this, we proceed in three steps described in detail in [Appendix C.2](#) and [C.3](#). In short, for each of the three cleaned datasets we first collapse all the articles for a particular day into one document, and then compute, using the estimated word distribution for each topic, the topic frequencies for this newly formed

metrics, but with factors (representing topics) constrained to live in the simplex and fed through a multinomial likelihood at the observation equation. [Blei \(2012\)](#) provides a nice layman introduction to topic modeling. More technical expositions of the LDA approach can be found in [Blei et al. \(2003\)](#) and [Griffiths and Steyvers \(2004\)](#).

⁶To further improve the reader’s understanding of what the different topics are (and are not), we investigate, in [Appendix B](#), how the topics relate to external texts freely available to the public.

US T0: Monetary Policy



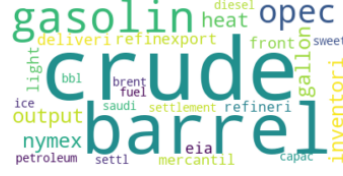
US T8: Strategy



US T55: Labor market



US T12: Petroleum



US T38: Stocks



US T78: Congress



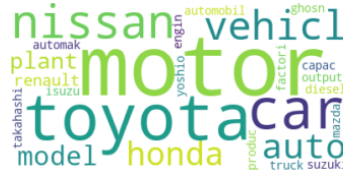
Japan T28: Outlook



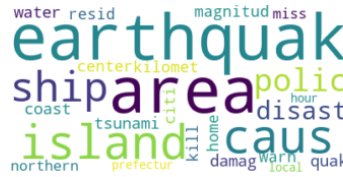
Japan T0: Russia



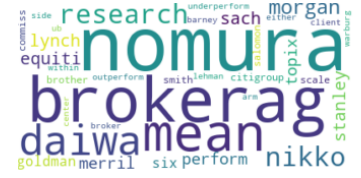
Japan T34: Motor



Japan T58: Natural disasters



Japan T33: Financial companies



Japan T46: Communication



Europe T5: Macroeconomics



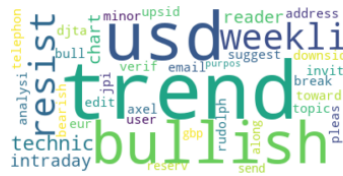
Europe T34: Trading data



Europe T48: Middle East



Europe T58: Investing



Europe T14: Fiscal policy



Europe T79: Health



Figure 2. Word clouds and topic categorization. For each word cloud the size of a word reflects the probability of this word occurring in the topic. Each word cloud only contains a subset of all the words in the topic distribution. Topic labels are subjectively given.

document. This yields a set of K daily time series. Then, for each day and topic, we find the article that is best explained by each topic, and from that identify the tone of the topic, i.e., whether or not the news is positive or negative. This is done using an external

word list and simple word counts, similar to in [Tetlock \(2007\)](#). The word list used here classifies positive/negative words as defined by the *Harvard IV-4 Psychological Dictionary*. For each day, the count procedure delivers a statistic containing the normalized difference between positive and negative words associated with a particular article. These statistics are then used to sign-adjust the topic frequencies computed in step one. Finally, we remove high frequency noise from each topic time series by using a 60-day (backward looking) moving average filter, and, as is common in factor model studies ([Stock and Watson \(2012\)](#)), standardize the resulting series. Figure 8, in Appendix A, illustrates the resulting series for the 18 word clouds presented in Figure 2.

Notice from the description above that also the tone adjustment procedure explicitly uses the output from the topic model. Still, the method used for identifying the tone of the news using dictionary based techniques is simple, and could potentially be improved upon with more sophisticated algorithms ([Pang et al. \(2002\)](#)). While leaving such endeavors for future research, [Thorsrud \(2016b\)](#) shows that working with topic frequencies without tone adjustment results in a loss of important information.

4 Business cycle narratives

To link the daily news topics time series to aggregate economic fluctuations, we estimate a coincident index of business cycles utilizing the joint informational content in quarterly output growth and the daily news narratives using a Dynamic Factor Model (DFM). This approach builds on conventional models proposed in, e.g., [Stock and Watson \(1988\)](#), [Mariano and Murasawa \(2003\)](#), [Aruoba et al. \(2009\)](#), and [Marcellino et al. \(2016\)](#), and has two important characteristics. First, since our best measure of aggregate economic fluctuations, changes in Gross Domestic Product (GDP), is observed at the quarterly frequency, the aggregation from higher to lower frequency variables is handled using a cumulator variable approach ([Harvey \(1990\)](#), [Banbura et al. \(2013\)](#)). Second, to summarize the informational content in the large panel of variables in a parsimonious manner, a factor modeling approach is implemented.

The novelty of the DFM used here is that we include daily news variables instead of hard economic statistics as observable variables (in addition to GDP), but also that the model allows for time-varying parameters with a latent threshold mechanism. This model property is motivated by our narrative definition (see Section 2), enforces dynamic sparsity, and has also proven to be important for both forecasting and more structural interpretation in other high-dimensional settings ([Zhou et al. \(2014\)](#), [Scott and Varian \(2013\)](#), [Thorsrud \(2016b,a\)](#)).

We obtain GDP statistics, measured in constant prices, for the US, Japan, and Europe

from the *Federal Reserve Bank of St. Louis* FRED database. The raw data is transformed into quarterly growth rates, and normalized. Then, a separate model is specified and estimated for each country. Following [Thorsrud \(2016b\)](#), and letting bold-font letters denote vectors and bold-font capital letters matrices, the DFM containing quarterly GDP growth and the daily news topic variables, can be written in a compact form as:

$$\mathbf{y}_t = \mathbf{Z}_t \mathbf{a}_t + \mathbf{e}_t \quad (1a)$$

$$\mathbf{a}_t = \mathbf{F}_t \mathbf{a}_{t-1} + \mathbf{R}_t \mathbf{\Sigma}_t \boldsymbol{\omega}_t \quad (1b)$$

$$\mathbf{e}_t = \mathbf{P} \mathbf{e}_{t-1} + \mathbf{u}_t \quad (1c)$$

with

$$\mathbf{y}_t = \begin{pmatrix} \mathbf{y}_t^{k_q} \\ \mathbf{y}_t^d \end{pmatrix} \quad \text{and} \quad \mathbf{a}_t = \begin{pmatrix} \mathbf{a}_t^{k_q} \\ a_t^d \end{pmatrix}$$

where t is the daily time index, k_q and d denote the quarterly and daily observation intervals, respectively, and the model has been written with simple autoregressive time series processes of order one for notational simplicity.⁷

Equation (1a) is the observation equation of the system. $\mathbf{y}_t^{k_q}$ and $\mathbf{y}_t^d$, are $N_q \times 1$ and $N_d \times 1$ vectors of quarterly and daily variables, respectively, with $N = N_q + N_d$. In this applications, $N_q = 1$ and $N_d = K = 80$. $\mathbf{Z}_t$ is a $N \times N_a$ matrix with dynamic factor loadings linking the variables in $\mathbf{y}_t$ to the latent dynamic factors in $\mathbf{a}_t$, and are described in greater detail below. The vector $\mathbf{e}_t$ contains the idiosyncratic errors. It is assumed that these evolve as independent AR(p) processes given by (1c), where $\mathbf{u}_t \sim i.i.d.N(0, \mathbf{U})$. Equation (1b) is the transition equation of the system. The common factors follow a VAR(h) process. $\boldsymbol{\omega}_t \sim i.i.d.N(0, \mathbf{I})$ and $\mathbf{\Sigma}_t$ is a diagonal matrix with $\mathbf{\Sigma}_t \mathbf{\Sigma}_t' = \mathbf{\Omega}_t$, allowing for stochastic volatility. The individual elements in $\mathbf{\Sigma}_t$ are assumed to follow random walk processes.⁸

The last element in $\mathbf{a}_t$, the scalar a_t^d , is interpreted as the latent common daily business cycle index. The other elements in $\mathbf{a}_t$, and in $\mathbf{F}_t$ and $\mathbf{R}_t$, contain cumulator variables used to handle the mixed-frequency property of the model. In the interest of brevity we describe the time aggregation procedure in Appendix D.7.

Dynamic sparsity is enforced on the system through the time-varying elements in $\mathbf{Z}_t$, which are modeled following the Latent Threshold Model (LTM) idea by [Nakajima and West \(2013\)](#). For one particular element in the $\mathbf{z}_t^d$ vector, $z_{i,t}$, the LTM structure can be written as:

$$z_{i,t} = z_{i,t}^* \varsigma_{i,t} \quad \varsigma_{i,t} = I(|z_{i,t}^*| \geq d_i) \quad (2)$$

⁷The model can easily be generalized to include variables of other frequencies as well (see [Thorsrud \(2016b\)](#) for details).

⁸While not explicitly discussed in this study, earlier studies show that allowing for stochastic volatility tend to improve the model performance in this type of DMFs (see, e.g., [Thorsrud \(2016a\)](#)).

where

$$z_{i,t}^* = z_{i,t-1}^* + w_{i,t} \quad (3)$$

with $w_{i,t} \sim i.i.d.N(0, \sigma_{i,w}^2)$, and $\mathbf{w}_t \sim i.i.d.N(0, \mathbf{W})$ where $\mathbf{W}$ is a diagonal matrix. In (2) $\varsigma_{i,t}$ is a zero one variable, whose value depends on the indicator function $I(|z_{i,t}^*| \geq d_i)$. If $|z_{i,t}^*|$ is above the threshold value d_i , then $\varsigma_{i,t} = 1$, otherwise $\varsigma_{i,t} = 0$.

The motivation for the LTM mechanism can easily be understood by an example. If the interpretation of narratives evolve and justify ongoing actions differently across time, or, if some narratives are more important in some periods than in others, a constant parameter model will fail. The researcher might simply conclude that a given narrative has no relationship with a_t^d , i.e., that $\mathbf{z}_i^d$ equals zero for all time periods, because, on average, periods with a positive $\mathbf{z}_i^d$ cancels with periods with a negative $\mathbf{z}_i^d$. The LTM mechanism potentially captures such cases in a consistent and transparent way.

A more detailed description of the time-varying DFM model, and estimation, is given in Appendix D. Here we note that the DFM is estimated by decomposing the problem of drawing from the joint posterior into a set of much simpler ones using MCMC simulations. Prior specifications are discussed in Appendix D.6.

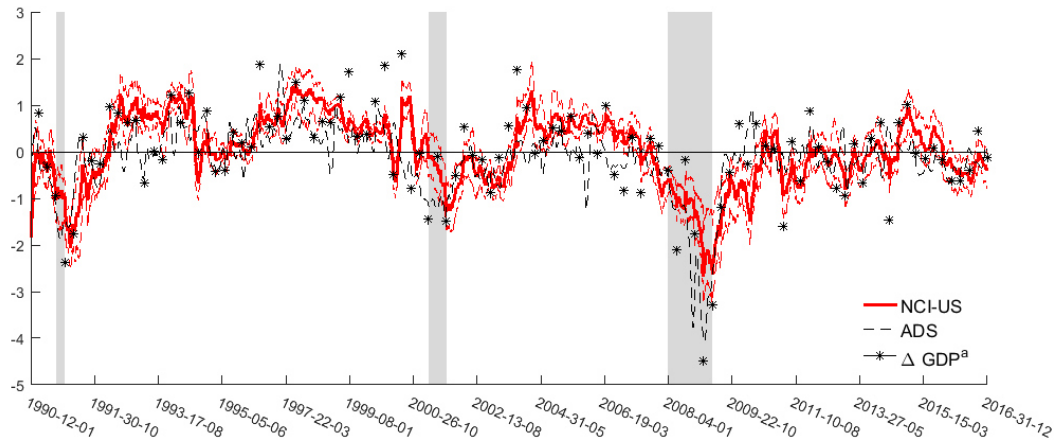
For all specifications we allow for one lag in the equation for the idiosyncratic errors ($p = 1$), and up to ten lags for the latent common business cycle index ($h = 10$). The (full) estimation sample ends 31 December 2016 for all three countries. Due to data availability, estimation starts in 12 January 1990, 29 June 1994, and 1 July 1995 for the US, Japanese, and European model, respectively. Finally, we globally identify the sign and size of the latent factor by restricting the factor loading for the first element among the N_d variables to equal 1 for all time periods. We choose the normalizing variables by looking at the simple correlation between linearly interpolated output growth and the daily news topics. Accordingly, for the US, Japan, and Europe we use the *Labor market*, *Outlook*, and *Macroeconomics*, news topics, respectively. Bai and Ng (2013) and Bai and Wang (2014) show that these restrictions uniquely identifies the factor and the loadings, but leaves the transition equation dynamics completely unrestricted.

4.1 The daily news-based coincident indexes

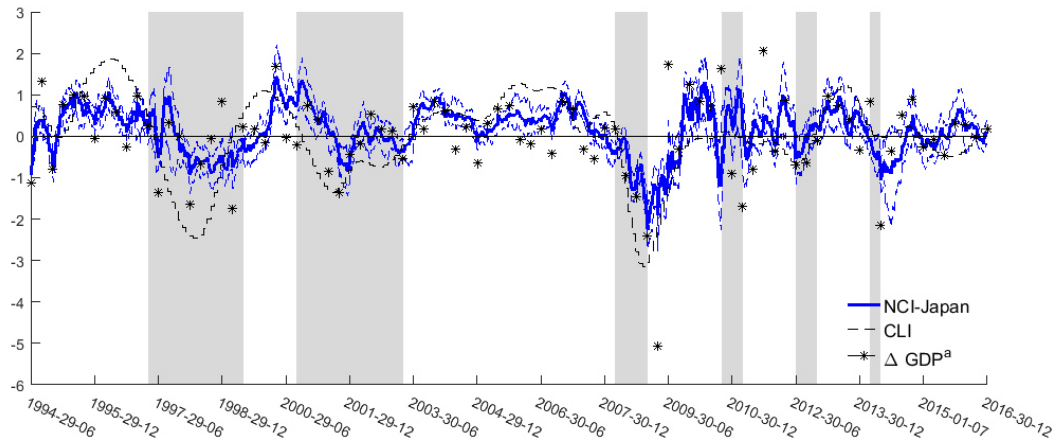
Figure 3 reports the estimated news-based coincident indexes for the US (*NCI-US*), Japan (*NCI-Japan*), and Europe (*NCI-Euro*). The gray shaded areas illustrate recession periods as defined by NBER (US), ECRI (Japan), and CEPR (euro area), while the black stars report observed quarterly GDP growth.⁹ In each graph we also report alternative existing state-of-the-art coincident index estimates. For the US, Japan, and Europe this is the

⁹NBER is the National Bureau of Economic Research, ERCI is the Economic Cycle Research Institute, while CEPR is the Centre for Economic Policy Research. Of these, only the chronologies provided by the

(a) NCI-US



(b) NCI-Japan



(c) NCI-Euro

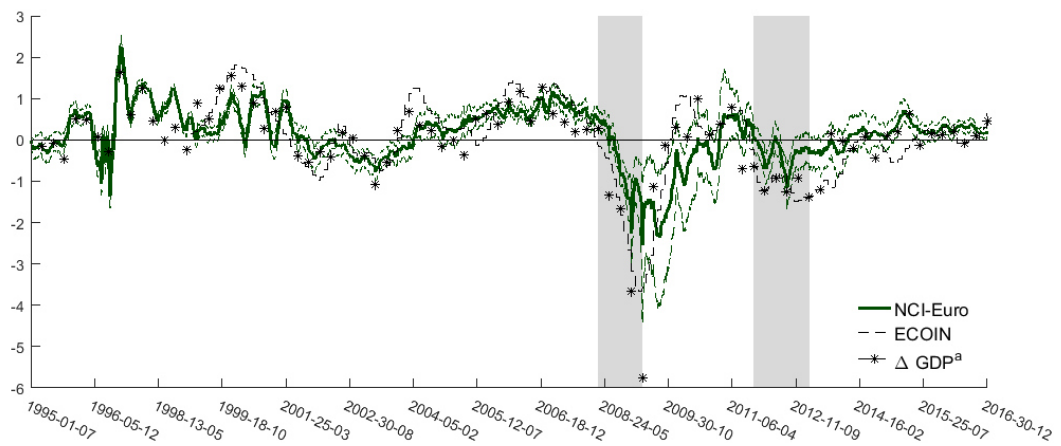


Figure 3. ΔGDP^a is standardized output growth. It is recorded at the end of each quarter. The colored solid line is the standardized (median) estimate of the daily business cycle index, while the dotted colored lines are the 68 percent probability bands. The gray shaded areas illustrate recession periods as defined by NBER (US), ERCI (Japan), and CEPR (euro area).

daily *ADS* index (Aruoba et al. (2009)), the monthly *CLI* index (Eurostat), and the NBER and CEPR are regarded as representing official business cycle dates. To the best of our knowledge, no official business cycle dating committees exists for Japan.

monthly *ECOIN* index (Altissimo et al. (2010)), respectively.

By simple visual inspection we observe that the estimated news-based indexes track the state of the economies very well, and that results for the US seem to be especially good. The financial crisis is common for all indexes, while the recession in the early 1990s is US specific. Likewise, the two long downturns in the late 1990s and early 2000s are specific for Japan, while the troubled times following the Great Recession are partly shared by both Japan and Europe. In relation to this, it is interesting to observe the substantial increase in uncertainty associated with *NCI-Euro* in the periods following the financial crisis.¹⁰

To formally evaluate the models we use classification tests. Like in Travis and Jordà (2011), and in the tradition of Burns and Mitchell (1946), we categorize aggregate economic activity into phases of expansions and contractions and evaluate the indexes' ability to classify such phases using Receiver Operating Characteristic (ROC) curves and area under the curve (AUROC) statistics. As measures of the unknown "truth", we use the business cycle chronologies illustrated in Figure 3, i.e., the business cycle phases defined by the NBER, ERCI, and CEPR. Since these chronologies are available at a daily frequency only for the US economy, daily classifications are obtained by assuming that the economies remain in the same phase on each day within the monthly classification periods for Japan and the euro area.

Focusing on the AUROC statistics, Table 1 summarizes the business cycle classification scores, while Figure 10 in Appendix A reports the associated ROC curves. As a perfect classifier receives an AUROC of 1, we observe from the table that the *NCI-US* index is tracking the official NBER business cycle chronology very well. Also the *NCI-Euro* index is doing a reasonably good job at classifying the phases of economic fluctuations. The worst performing index, in terms of AUROC, is *NCI-Japan*, which receives a score of 0.76. Still, this is far better than random guessing, which would give an AUROC of 0.5.

To put the performance of the news-based indexes into perspective, we also evaluate the classification performance of the alternative state-of-the-art coincident indexes illustrated in Figure 3. Of these, only the *ADS* index is available on a daily frequency. For the monthly *CLI* and *ECOIN* indexes we construct daily analogs by assuming that every day within a month equals the observed monthly value. Again, Table 1 summarizes the results. In all three countries the existing indexes perform slightly better than the news-based indexes. However, the differences are not large, and at most 12 percent, for the euro area. In addition, the news-based indexes are available at a daily frequency, which

¹⁰The time-varying changes in the variance of the *NCI* errors are illustrated in Figure 9 in Appendix A. Unexpectedly, all models pick up a substantially higher variance during the financial crisis episode than in other parts of the sample. Convergence statistics indicating that the MCMC algorithm has reached the ergodic distribution are discussed in Appendix F.

Table 1. Receiver Operating Characteristics and area under the curve (AUROC) statistics. By definition the AUROC can not exceed 1, perfect classification, or be lower than 0.5. We compute the AUROC score non-parametrically using the algorithm described in [Travis and Jordà \(2011\)](#).

	<i>NCI-US</i>	ADS	<i>NCI-Japan</i>	CLI	<i>NCI-Euro</i>	ECOIN
AUROC	0.946	0.996	0.760	0.790	0.853	0.969

the alternative indexes typically are not.

In sum, these results illustrate how informative the news-based approach is in terms of capturing economic fluctuations. For countries where high-frequency hard economic variables are not easily available, the news-based approach offers a valuable alternative.¹¹ Moreover, in contrast to existing coincident indexes, the news-based approach gives the researcher, or index user, potential knowledge about the narratives important for understanding economic fluctuations. An issue we now turn to.

4.2 Business cycle decompositions

In this section we investigate “the epidemiology of narratives relevant to economic fluctuations” ([Shiller \(2017\)](#)). We do so by utilizing an attractive feature of the DFM modeling framework, namely that the state evolution of the model (the daily business cycle index(es)) can be decomposed into news surprises driven by the developments in the observable variables (the news topics). Technically, this is done using Kalman Filter iterations and decomposing the state evolution at each updating step into news contributions using the Kalman Gain (see [Appendix E](#)), and the recursive nature of the filter. Following [Koopman and Harvey \(2003\)](#), let:

$$\mathbf{a}_{t|t} = \mathbf{a}_{t|t-1} + \mathbf{K}_t \mathbf{v}_t \quad (4)$$

be the standard Kalman filter equation for updating the latent state estimate $\mathbf{a}_t$ given knowledge of the Kalman Gain matrix $\mathbf{K}_t$, with:

$$\begin{aligned} \mathbf{a}_{t|t-1} &= \mathbf{F}_t \mathbf{a}_{t-1|t-1} \\ \mathbf{v}_t &= \mathbf{y}_t - \mathbf{Z}_t \mathbf{F}_t \mathbf{a}_{t-1|t-1} \end{aligned} \quad (5)$$

Now, plugging (5) into (4) one obtains:

$$\begin{aligned} \mathbf{a}_{t|t} &= \mathbf{F}_t \mathbf{a}_{t-1|t-1} + \mathbf{K}_t (\mathbf{y}_t - \mathbf{Z}_t \mathbf{F}_t \mathbf{a}_{t-1|t-1}) \\ &= (\mathbf{I} - \mathbf{K}_t \mathbf{Z}_t) \mathbf{F}_t \mathbf{a}_{t-1|t-1} + \mathbf{K}_t \mathbf{y}_t \end{aligned} \quad (6)$$

¹¹Although the DFM model, with the LTM mechanism, is built to filter out uninformative data, it is very likely that a more elaborate data (pre)selection procedure could improve the results further. High frequency (hard) economic indicators can also be included into the model alongside the news topic variables. We leave such attempts for future research.

Table 2. Top 10 news topic (surprises). The ranking is based on sorting the output from equation (7) in descending order.

<i>NCI-US</i>	<i>NCI-Japan</i>	<i>NCI-Euro</i>
Labor market	Outlook	Macroeconomics
Stocks	Motor	Middle East
Monetary policy	Financial companies	Trading data
Clients	Fed	Fiscal policy
Congress	Russia	Bonds
Regulations	Stock listings	Credit rating
Strategy	Market commentary	Nordic countries
Petroleum	Natural disasters	Australia
Education	Communication	Public safety
Market performance	Car technology	Investing

which can be inverted to obtain the moving average representation of the unobserved states as a function of the observed variables. Or, in other words, how the model interprets surprising news fluctuations when updating the state estimates.

Defining $w_{i,t} = K_{i,t}v_{i,t}$ as the weighted forecast error contribution from topic i at time t , and:

$$w_i = \frac{1}{T} \sum_{t=1}^T (w_{i,t})^2 \quad (7)$$

as the mean squared error, Table 2 reports the 10 most influential news topics on average across the sample. In general, news surprises about macro economic developments (e.g., *Labor market*, *Macroeconomics* and *Outlook*), the financial market (e.g., *Stocks* and *Trading data*), and (geo-)politics (e.g., *Monetary policy*, *Fiscal policy*, *Congress*, *Middle East*, and *Russia*) are important in all three countries. Still, constructing a story based on words drawn from the topic distributions summarized in the three columns in Table 2 would clearly result in three different narratives. For example, a grand narrative about Japan would be much more likely to contain topics related to the motor and car industry, and natural disasters, than a story for the US or euro area. Likewise, for a US-specific story, topics related to *Petroleum* and *Regulations* are likely much more prominent than in any of the other two countries.

Table 3, for the US, and Tables 13 and 14 in Appendix A, for Japan and the euro area, list the most influential narratives across six different sub-samples, as well the first sentences of particularly representative news articles during these time periods. While some of the same news topics tend to top the lists in every period, we observe a relatively large variation in the ranking of the other narratives. For example, during the period 1999-2002, topics associated with *Internet* and *Persuasion* are in the top of the list for the US,

Table 3. Top five news topics across sub-samples for the *NCI-US* index. The *Example narratives* are found by querying the corpus for news articles where the five news topics listed in column two combined receive a high weight. Only the first sentences of each story are included in the table. The date of publication is printed in parenthesis.

	Top 5 news topics	Story example
1995 - 1999	Labor market Europe Market perfor. East Asia Petroleum	(1996-04-24) <i>Western Germany's consumer price index (CPI) is estimated to have risen a preliminary 0.2% in April from February and 1.5% from a year ago, a survey conducted by AP-Dow Jones shows... Economists concurred that the expected increase in the price index is largely due to an increase in energy prices...</i>
1999 - 2002	Labor market Education Design Internet Persuasion	(2000-05-22) <i>So you've started a successful company before your 30th birthday. Big deal. Navin Chaddha has co-founded five. What's more, the 29-year-old electrical engineer has assisted and even invested hundreds of thousands of his own dollars in at least eight other start-ups...</i>
2002 - 2006	Stocks Labor market Events Terrorism Strategy	(2003-09-05) <i>Look past the ongoing sabotage and strife in Iraq and you will see that the Bush administration is eager to pull off the most ambitious economic reform in a Middle Eastern country since the dissolution of the Ottoman Empire... The administration wants to promote free trade for the entire gamut of Arab countries...</i>
2006 - 2009	Labor market Regulations Congress Natural gas Strategy	(2008-09-17) <i>Congressional auditors are questioning whether the Interior Department is collecting all the royalties energy companies owe for petroleum developed on federal property... Last year, the MMS collected more than \$11.4 billion in oil, natural-gas and other mineral royalties... Congress this week is debating proposals to allow more offshore oil drilling...</i>
2009 - 2013	Labor market Clients Elections Sports Congress	(2010-03-03) <i>When it comes to talking about what is holding back the economy, politicians in Washington should look in the mirror. Inaction and infighting on the government level have resulted in a loss in confidence among consumers and business owners that their elected officials are doing the right thing when it comes to healing the economy or bringing down unemployment...</i>
2013 - 2016	Labor market Monetary policy Documentation Clients Design	(2013-05-16) <i>Even though inflation measures have fallen sharply in recent months, Federal Reserve officials aren't ringing alarm bells about it as they have done in the past. Fed officials have said they take comfort that the public's expectation of future inflation, as registered in surveys of households and bond markets, has remained stable...</i>

whereas the topic *Terrorism* enters the list during the 2002-2006 period. Likewise, the *Terrorism* narrative enters the top five list during the 2013-2016 period in the euro area together with the *Monetary policy* topic. Interestingly, and something we will come back

to, the narrative focus on monetary policy is also shared by the US and Japan during this time period. The news article excerpts reported in the tables illustrate how the discovered topic structure in the corpus, together with the DFM decomposition, provides meaningful mappings. It is, for example, easy to argue that the excerpts for the US are about at least *Europe*, *Petroleum*, and *Market performance* (1995-1999), *Regulations*, *Congress*, and *Natural gas* (2006-2009), and *Labor market*, and *Congress* (2009-2013).

Figure 4 provides an illustration of how news surprises in the US affect the *NCI-US* estimates over time, at a daily frequency. Two distinct results stand out. First, the timing of when specific topics become important, either positively or negatively, resonates well with the conventional narrative held about economic developments the last two decades. At the risk of cherry picking, we give some examples: Prior to, and going into the 2001 recession, surprising news related to the *Internet*, *Design*, *Education*, *M&A*, and *Volatility* topics pulled the coincident index upwards, while narratives related to *Labor markets*, *Bankruptcies*, and *Automobiles* pulled the coincident index downwards. Thus, interpreted through the lenses of the model proposed here, the burst of the dot-com bubble is well identified, but the news topic developments directly related to the grand dot-com narrative was not as bad as the model expected. Conversely, news topic developments related more towards the general economic conditions came in worse than predicted. The story related to the financial crisis in 2007/2008 is of a somewhat different type. Now surprising negative movements in topics as *Strategy*, *Bonds*, and *Regulations*, stand out. Lastly, turning to the slow recovery period following the financial crisis, we observe that unexpected news about *Congress*, *Economic crisis*, *Funding*, *Environment*, and *Commodities* contributed negatively to growth, while topics related to *Labor market*, *Sports*, *Commentary*, *Natural gas*, and *Elections* helped pull the index upwards.

Second, the degree of sparsity enforced on the factor loading space changes considerably across time. For example, during the 1990s few factor loadings have a high probability of being zero. In the period following the financial crisis, however, the degree of sparsity is much larger, with only a few time-varying factor loadings being larger (in absolute value) than their respective threshold. It is also interesting to see how the degree of sparsity seems to increase around recession periods. That is, when times are bad, our results indicate that the set of narratives relevant for economic fluctuations is smaller. Interestingly, this finding is very much in line with theory models explaining how news coverage becomes more homogeneous around major events, and thereby increasing the correlation among economic agents' actions (Nimark and Pitschner (2016)). Thus, in relation to narratives, booms are broad-based while busts are not.

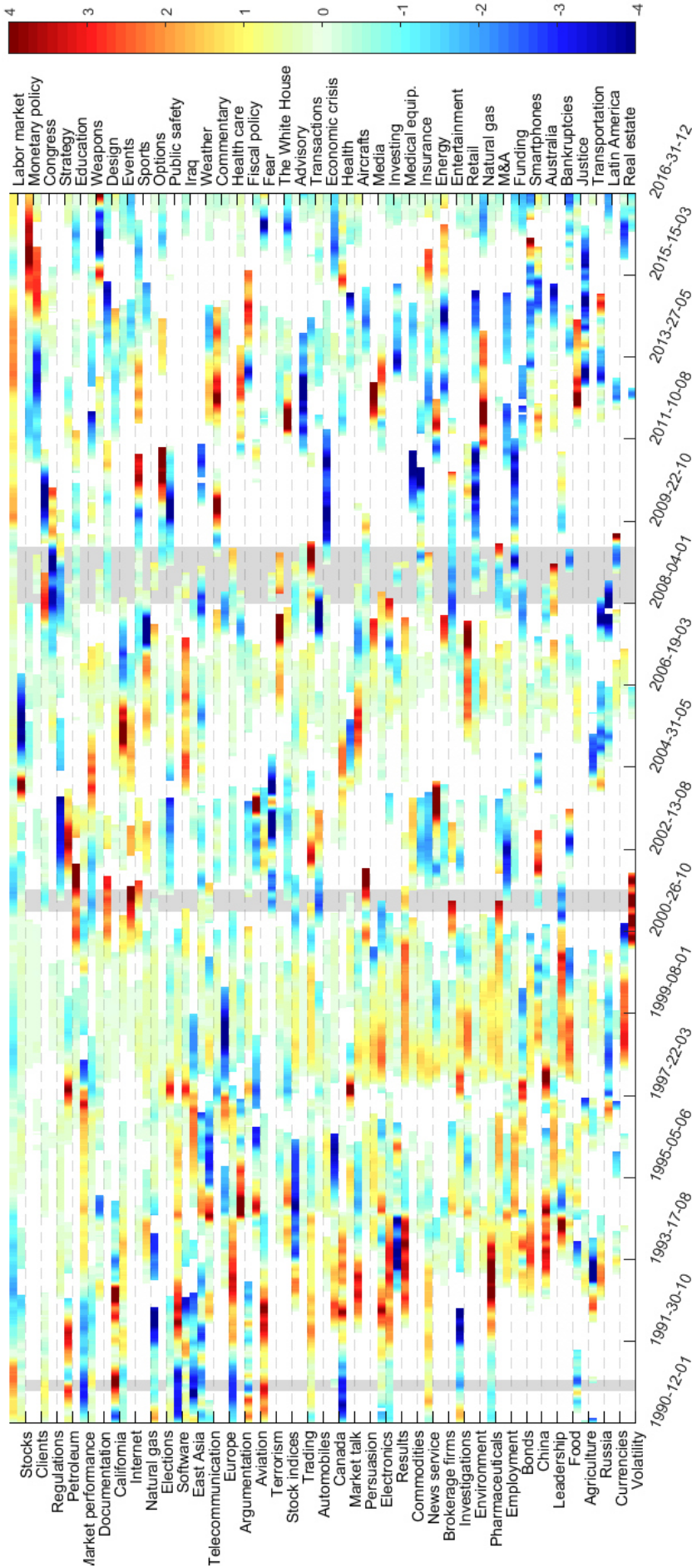


Figure 4. US news topics and their contribution to coincident index estimates across time. The reported decompositions are based on running the Kalman Filter using the posterior median estimates of the hyper-parameters and the time-varying factor loadings (at each time t). In the interest of readability, the topic names are reported on two y-axes with two-step increments. For example, the *Labor market* topic is associated with the first row (from above) in the figure, while the *Stocks* topic is associated with the second row (from above). White areas illustrate the time-varying sparsity patterns. Recession periods, defined by NBER, are illustrated using gray shading.

Figures similar to 4 are reported for Japan and the euro area in Figures 11 and 12 in Appendix A. Instead of going into the details, we highlight that we see clear sparsity patterns around recession periods, like in the US. In Europe, for example, *Credit rating*, *Bonds*, *Investing*, *Outlook*, and *Funding* are almost the only news topics contributing to explaining the negative developments in the euro-area business cycle index during, and following, the financial crisis. Similarly, in Japan narratives related to *Electronics*, *Retail*, *Income*, and *Growth* contributed especially negatively during 2009, while the period between 2010 and 2011 is partly dominated by negative news topic surprises attributed to *Politics* and *US politics*.

Finally, although most topics are easily interpretable and provide information about what is important for the current state of the economy, some topics either have labels that are less informative, or reflect surprising categories. From the US-based decompositions, in Figure 4, examples are the *Sports*, *Entertainment*, and *Food* topics. That said, such exotic or less informative named topics, are the exception rather than the rule. It is also the case that a news article is a mixture of topics. To the extent that different topics, meaningful or not from an economic point of view, stand close to each other in the decomposition of the corpus they might covary and therefore both add value in terms of reflecting the current state of the economy.

We conclude that the decompositions of the business cycles into narrative contributions tell a story about economic fluctuations reasonably in line with historical experience. This should not be too surprising, given that the narratives we know are the ones we have been served, partly through the media. What is perhaps more surprising is that it is quantified so well. The finding about narrative sparsity around recessions is novel, and some of the influential news topics clearly represent (economic) concepts or events that would have been very difficult, if not impossible, to capture using conventional economic data.

4.3 Going viral?

Shiller (2017) argues, but does not quantify, that “narratives “go viral” and spread far, even worldwide, with economic impact”. Accordingly, a reasonable testable hypothesis is that there exists a significant relationship between how important similar news topics are in explaining business cycle developments across countries and economic fluctuations, at least periodically. We investigate this hypothesis by first constructing statistics measuring how similar the news topics are across countries. Then, we weight these similarity measures with how important the news topics are in explaining business cycle developments and derive what we label virality indexes. These indexes give a quantitative measure of the degree to which (similar) narratives relevant for growth go viral. Finally, we exploit the high frequency nature of our data, and investigate if there is any significant relationship

between the virality indexes and economic fluctuations across countries.

To measure topic similarity across countries, we use the Jensen-Shannon divergence (JSD). This is a method for measuring the similarity between two probability distributions. The JSD is based on the Kullback-Leibler divergence, but it is symmetric, always a finite value, and bounded between 0 and 1. Formally, for two discrete probability distributions P and Q :

$$JSD(P||Q) = \frac{1}{2}D(P||M) + \frac{1}{2}D(Q||M) \quad (8)$$

where $M = \frac{1}{2}(P + Q)$, and $D(P||M)$ is the Kullback-Leibler divergence:

$$D(P||M) = \sum_i P_i \log_2 \frac{P_i}{M_i} \quad (9)$$

Here, with reference to Section 3.2, and Table 1, P and Q are two word distributions (Φ_k) associated with two different topics. Treating the US economy as the common “numeraire”, we compute the $JSD(P||Q)$ for all combinations of topics in the US and either Japan or Europe. This results in two $K \times K$ matrices, one for each country pair, with JSD scores. Table 12, in Appendix A, reports the topic combinations with the lowest JSD score (most similar), and shows that the mappings make sense intuitively. For example, the US topics we have labeled *Fiscal policy*, *Funding*, and, *Telecommunication*, have gotten the same labels in both Japan and Europe, while the US topic *Monetary policy* has gotten the label *Fed/BoJ* and *Fed* in the Europe and Japan, respectively. In some cases, however, there are larger, less intuitive, discrepancies. An example is the US-based topic labeled *Canada* by us, which according to the JSD score is most similar to the European and Japanese topics *Outlook* and *Fiscal policy*.

The virality index $VIR_t^{s,US}$ between country s and the US is constructed as follows:

$$VIR_t^{s,US} = \sum_{i=1}^{80} \sum_{j=1}^{80} \left[\frac{\tilde{w}_{t,i}^s \tilde{w}_{t,j}^{US}}{(c + JSD_{i,j}^{s,US})} \right] \quad s = \{Japan, Euro\} \quad (10)$$

Here, $\tilde{w}_{t,i}^s = w_{i,t} / \sum_i^K w_{i,t}$, with $w_{i,t}$ defined in Section 4.2, i.e., the normalized weight given to topic i in explaining the movements in the business cycle index in country s at time t , while the $JSD_{i,j}$ term defines how similar topic i in country s is to topic j in the US. c is a small constant ensuring that we do not divide the expression by 0, which is the lower limit of the VIR indexes.

Figure 5 reports the two virality indexes. On average, the indexes fluctuate mildly. However, at times the indexes spike, and some narratives go viral and become an epidemic. This pattern is especially pronounced following the financial crisis in 2008, when the frequency, duration, and magnitude of the spikes all increase significantly relative to the periods before. More formally, using a peak-finding algorithm to compute the number

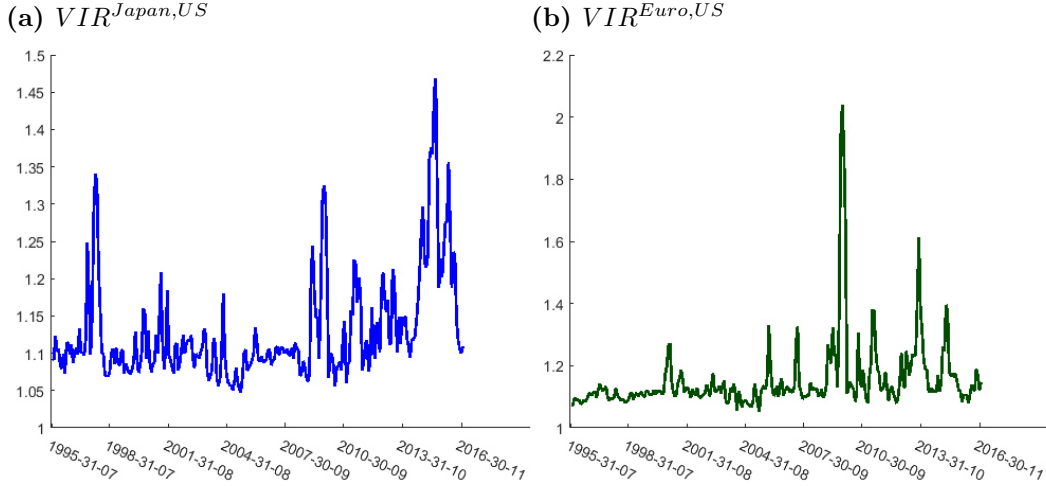


Figure 5. Virality indexes for the US-Japan and US-Europe economies. In the interest of visual clarity, the indexes are plotted on a monthly frequency, where aggregation from daily to monthly frequency is obtained by a simple mean.

of peaks, and their duration, we identify only two peaks prior to 2008, see Figure 13 in Appendix A. This is in the late 1997 for the $VIR^{Japan,US}$ index, and in early 2000 for the $VIR^{Euro,US}$ index. The length of these episodes are roughly 3 and 6 months. In contrast, in the periods following 2008, we identify in total 11 epidemics with durations up to 8 months.¹² The average duration of the epidemics are estimated to be around 5 and 4.5 months for the $VIR^{Japan,US}$ and $VIR^{Euro,US}$ indexes, respectively, where events happening late in the sample tend to pull these averages up.

Borrowing from Shiller (2017) and the spread of disease literature and the benchmark SIR model of Kermack and McKendrick (1927), our results indicate that the contagion rate (co) to recovery rate (re) ratio has increased over time. That is, (narrative) epidemics in the post 2008 period are more severe than in previous periods. Many different explanations can rationalize this finding. It is for example easy to argue that the introduction of internet and social media likely have increased both co and re (Zhao et al. (2013)). However, in terms of Figure 5, it seems strange that this should have happened exactly in the mids of the financial crisis in 2008, suggesting instead that the epidemics observed during and after 2008 might be of a very different type than those encountered during the 1990s and early 2000s.

In Figure 13, in Appendix A, we also report the topic mappings contributing the most to the VIR estimates during the epidemic periods discussed above. Three broad findings stand out. First, epidemics are mostly associated with the US *Labor market* topic. In

¹²We have also tried defining periods of virality using a generalized version of the sup augmented Dickey-Fuller test (Phillips et al. (2015)). However, this test has low power in terms of correctly classifying spikes/bubbles when the duration of each is small relative to the total sample size. As this is the case here, the number of periods defined as explosive are far fewer than suggested by Figure 5.

Table 4. Epidemics and economic fluctuations. For each month in the sample we compute the mean and standard deviation of the three news-based coincident indexes, as well as their correlation with the *NCI-US* index, using the daily observations. Contagion periods (Cont.) are defined using the timing and durations implied by the results in Figure 13, in Appendix A. Periods of no contagion are defined as normal times (Norm). Significant differences in the moments (Diff) are tested using the Welch’s t-test. The superscripts ***, **, and * denote the 1% , 5%, and 10% significance level, respectively.

	US			Japan			Europe		
	Cont.	Norm	Diff	Cont.	Norm	Diff	Cont.	Norm	Diff
E(X)	-0.30	0.06	-0.36**	-0.22	0.21	-0.43**	-0.01	0.14	-0.16
STD(X)	0.07	0.07	-0.00	0.11	0.08	0.03**	0.07	0.07	0.00
COV(X, US)				0.01	-0.03	0.04	-0.04	0.15	-0.19

almost all episodes this topic features as a central component in the explaining the spikes in the VIR indexes. Second, there are three exceptions to this first point, namely the spike in the $VIR^{Euro,US}$ index in 2000, and the spikes in the $VIR^{Japan,US}$ index in 2014 and 2015. The former is undoubtedly related to the burst of the dot-com bubble, while the two latter are associated with the US *Monetary policy* topic. Third, the diversity of topics needed to explain a sizable share of the epidemic episodes varies considerable across time. During the spike in the $VIR^{Euro,US}$ index in September 2009, only one topic mapping is needed to explain up to 40 percent of the index. In contrast, during the September 2013 epidemic in the same index, 13 topic mappings are needed. Thus, some epidemic episodes have a “sharp” narrative interpretation, while others are more complex. Based on the topic contributions, and the timing, we can for example conjecture that the 2009 episodes are related to the Great Recession, while the 2011 episodes are related to the massive earthquake that hit Japan this year, sparking well known global concerns about both finance, trade, and energy related topics. We do not find, however, any relationship between the estimated duration of the epidemics, and the number of topic mappings needed to explain a sizable share of the VIR indexes during such episodes.

The estimated timing of the VIR epidemics suggest that they are associate with bad events, and thus potential negative economic developments. The results reported in Table 4 confirms this impression. Higher values of the VIR indexes are associated with lower growth rates than in “normal” times in all three countries, and significantly so in the US and Japan. On the other hand, we do not find any significant differences in the covariances between the country pairs during periods of epidemics relative to normal times. If anything, it becomes lower between the US and Europe. To the extent that increases in the $VIR^{Euro,US}$ index are considered as some type of common shock(s) to the international business cycle (Kose et al. (2003), Stock and Watson (2005)), this means that

their (short-term) propagation differ across countries, potentially leading to divergence, as opposed to convergence, of international business cycles (Mumtaz et al. (2011) Kose et al. (2012)).

To summarize, the preceding analysis has shown that narratives do “go viral” and spread worldwide, as argued by Shiller (2017), but mostly so in times of trouble. The narratives contributing the most to the epidemic episodes tend to be associated with US-based macroeconomic developments and (partly) monetary policy.

4.4 Behind the news

No causal inference is sought, or can be inferred, from the preceding analysis. Here, we take one step towards a more structural understanding of information diffusion. In particular, we ask how narratives independently spread between economic regions, and whether news topics that are important for describing, e.g., the US business cycle, have predictive power for narratives in Japan and Europe, or vice versa.

We answer these questions by building on the well known Granger causality concept (Granger (1969)). A variable is said to Granger cause another, if the first series contains additional information for predicting the future values of the second series, beyond the information in the past values of this second series. While originally formulated in a low dimensional setting, recent work trying to infer causal relationships among components of biological systems has extended this reasoning to high dimensional problems through the usage of “Graphical Granger causality” modeling (Lozano et al. (2009), Shojaie and Michailidis (2010)). These methods offer efficiency gains over more standard (pairwise) Granger causality tests because of the usage of regression methods with variable selection and regularization (Arnold et al. (2007)), i.e., Lasso and its variants, and are tailored for high dimensional problems, as here.

Let $\mathbf{y}^j = (y_1^j, \dots, y_T^j)'$ be a $T \times 1$ response variable j , and $\mathbf{X} = [\mathbf{X}^1, \dots, \mathbf{X}^J]$ be the predictor matrix for $j = 1, \dots, J$ groups of covariates (including $\mathbf{y}$). Each matrix $\mathbf{X}^j = [L^1 \mathbf{x}^j, \dots, L^h \mathbf{x}^j]$, where $\mathbf{x}^j = (x_1^j, \dots, x_T^j)'$, L is the lag operator and h is the maximum number of lags. Then, we answer the question posted above, i.e., how narratives independently spread between economic regions, by estimating the group Lasso of Yuan and Lin (2006):

$$\hat{\beta}^j(\lambda) = \operatorname{argmin}_{\beta} \|\mathbf{y}^j - \mathbf{X}\beta\|^2 + \lambda \sum_{j=1}^J \|\beta_{G_j}\|_2 \quad (11)$$

for each $j = 1, \dots, J$. $\beta_{G_j} = \{\beta_k; k \in G_j\}$ and G_j denotes the set of group indexes. In our case each group is of equal length, and correspond to all the lagged variables belonging to one group. λ is the Lasso regularization parameter that shrinks or sets some of the groups (coefficients) to 0. Thus, the group Lasso is faithful to the original (pairwise) Granger

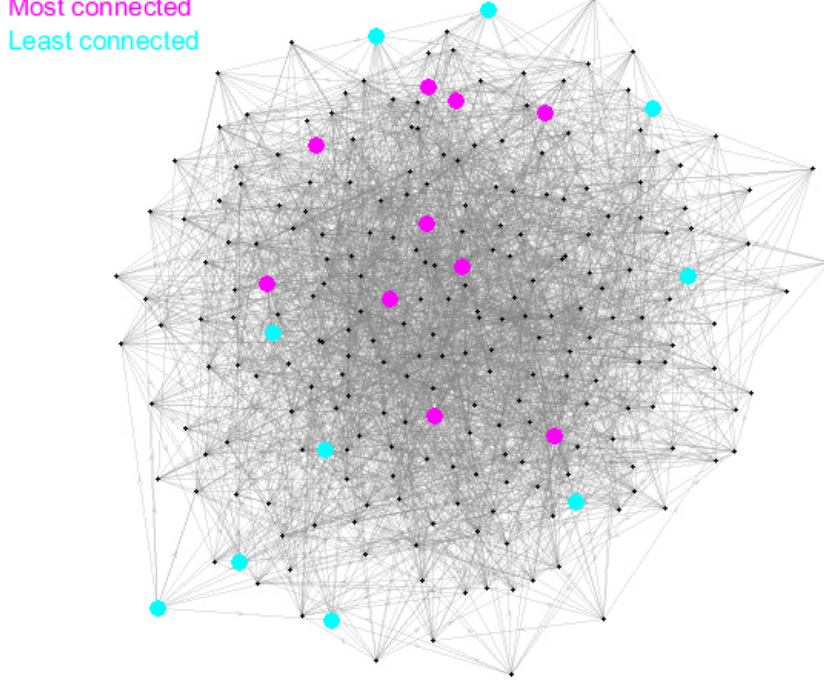


Figure 6. A network graph of the graphical Granger causality results. Each node is a narrative (news topic time series). In the interest of visual clarity, their name is not reported. The (gray) edges connecting the nodes are directed, and illustrate the direction of predictability across narratives. The highlighted nodes are those that are estimated to be the most (least) central narratives in the graph, see Table 6.

causality concept, where $\mathbf{x}^{i \neq j}$ is said to Granger cause $\mathbf{y}^j$ only if the entire lagged series $\mathbf{X}^{i \neq j}$ provides additional information for the prediction of $\mathbf{y}^j$.

We have $J = K \times 3 = 80 \times 3 = 240$ individual news topic time series, or groups. Before estimation, to reduce noise, the individual news topic time series are aggregated to monthly values, and the predictor matrix is standardized to make estimation scale invariant. We consider up to a half-year of lags, with $h = 3$. As $T \ll (J \times h)$, a standard regression framework is infeasible, while the Lasso applies because of the regularization term. For each j , we set λ_{max}^j such that it gives the largest non-null model. The group Lasso solution path is then computed by evaluating on 100 equally spaced λ 's between 0 and λ_{max}^j . The optimal λ_{opt}^j is chosen based on the BIC, as in [Lozano et al. \(2009\)](#).

We focus on cross-country spillovers, and say that a topic i in country s_1 Granger causes topic j in country s_2 when $\hat{\beta}_{G_{i \neq j}}^j(\lambda_{opt}^j) \neq 0$. More generally, [Shojaie and Michailidis \(2010\)](#) show how the output from procedure described above admits a graphical interpretation. In particular, we can construct the adjacency matrix of a directed acyclic graph (DAG) by stacking the estimated $J \times 1$ coefficient vectors $\hat{\beta}^j(\lambda_{opt}^j)$, for $j = 1, \dots, J$, into a $J \times J$ matrix A , whose (i, j) -th entry indicates whether there is an edge (and its weight) between nodes i and j . Below, to simplify the interpretation, we do not count relationships where there is a two-way predictive relationship, and set elements in A where both the (i, j) -th and (j, i) -th are non-zero to 0.

Table 5. Graphical Granger causality. For each country, the *Tot.* columns report the number of outgoing edges in percent of total potential connections. The remaining columns decomposes this fraction into cross-country contributions..

US			Japan			Europe		
Tot.	Japan	Europe	Tot.	US	Europe	Tot.	US	Japan
9.92	49.06	50.94	4.91	48.73	51.27	6.66	36.38	63.62

The network graph in Figure 6 illustrates the complexity of the problem, and shows that the interconnectedness of narratives across the US, Japan, and Europe is large. Still, given the large number of potential connections, the density of the graph is rather small, and estimated to be approximately 5 percent.¹³ The statistics reported in Table 5 break the graph density into country specific contributions. Out of 12800 potential connections, US-specific narratives dominate, and Granger cause roughly 10 percent of the foreign news topics. The direction of predictability is divided equally towards topics in Japan and Europe. The importance of Japan is only half that of the US, while narratives classified as being euro area-specific Granger cause roughly 7 percent of the foreign topics. However, in contrast to the results for the US- and Japan-specific news topics, the direction of the European-specific predictability is clearly tilted towards Japan.

As these results are new, they are hard to compare to existing knowledge. Still, the US-based dominance is well in line with common perception, and adds to the evidence about the US’s role in the global economy more broadly (Kose et al. (2017)).¹⁴

To gain knowledge of the narratives’ importance in the graphical Granger causality network, and relate this importance to the narratives’ importance for economic fluctuations, we compute a measure of the graph node’s centrality using the much applied “betweenness” measure (Freeman (1977)). This centrality metric measures how often each graph node appears on a shortest path between two nodes in the graph, and is computed as:

$$c(u) = \sum_{i,j \neq u} \frac{n_{ij}(u)}{N_{ij}} \quad (12)$$

where $n_{ij}(u)$ is the number of shortest paths from i to j that pass through node u , and N_{ij} is the total number of shortest paths from i to j . In addition, a cost, equaling $1/\tilde{w}_i^s$,

¹³The density of the graph is computed as the number of non-zero elements in the adjacency matrix A relative to the number of total elements.

¹⁴In unreported results we confirm that these findings hold when partitioning the sample into three equally sizes sub-samples, and re-estimating the graphical Granger causality graph for each. Relatedly, Table 8, in Appendix A, shows that among the daily business cycle indexes themselves, neither the *NCI-Japan* nor the *NCI-Japan* index Granger cause the *NCI-US* index, while the *NCI-US* index Granger causes at least the *NCI-Euro* index.

Table 6. Topic centrality. The centrality ranking is computed using the weighted “betweenness” measure of [Freeman \(1977\)](#). The *In degree* and *Out degree* counts reflect how many series that predict the listed topics, and how many topics the listed topics themselves predict, respectively.

Most central			Least central		
Name	In degree	Out degree	Name	In degree	Out degree
Euro T5-Macroeconomics	8	20	Japan T55-Intervention	19	10
Euro T48-Middle East	8	11	Euro T66-Justice	9	11
US T30-Regulations	9	12	Japan T19-Months	2	11
Japan T6-Fed	13	11	US T25-Clients	0	13
US T55-Labor market	10	14	US T28-Software	0	20
Euro T14-Fiscal policy	17	10	US T38-Stocks	0	17
US T16-Market performance	6	18	US T57-Australia	2	17
Japan T62-Car technology	14	10	Euro T52-Credit rating	0	15
Japan T3-Aviation	21	11	US T65-Bankruptcies	0	21
Japan T58-Natural disasters	28	3	US T74-Commodities	0	17

is assigned to each edge in the graph, where $\tilde{w}_i^s$ was defined in Section 4.3 as the average normalized weight given to topic i in explaining the movements in the business cycle index in country s . Thus, when computing the shortest path between two nodes in the graph, we rather traverse across edges which are important for explaining aggregate economic fluctuations.

Table 6 reports the 10 most and least important narratives according to (12). The news topics that are found to be important for explaining the economic fluctuations in the US, Japan, and Europe (confer Table 2), are also among the most important narratives in the graphical granger causality graph. The *Macroeconomics* topic in Europe, for example, is at the top of the list, and has an in and out degree in the network of 8 and 20, respectively. Conversely, at the bottom of the list we find the US *Commodities* topic. This news topic times series is not predicted by any of the other narratives, and therefore has a very low $c(u)$ ranking. Still, even though many of the least central news topics have low in degree, many of them have a relatively high out degree. The colored nodes in Figure 6 illustrate this, where the narratives with a low $c(u)$ score tend to be found far out in the network graph, while more connected news topics tend to be found closer to the center of the graph.

Figures 14 and 15, in Appendix A, provide examples of how two of the most and least central narratives in the network graph are connected to other topics. The figure is constructed as a subgraph of Figure 6. As the figures illustrate, the narratives *Macroeconomics* and *Commodities* tend to predict narratives of a similar type in other countries. For the *Middle East* and *Bankruptcies* topics, however, the predictive relationships are

more diverse.

Lastly, it is worth noticing before turning to the next section that among the 10 least central narratives in Table 6, we find 6 US news topics. Of these, both the *Stocks* and *Clients* are also among the 10 most important in terms of describing the US business cycle, confer Table 2 in Section 4.2.

4.5 News or noise?

The literature we speak to is divided in its view on whether narratives contain fundamental economic information, or just noise and sentiment. One branch of the literature can be associated with the news-driven business cycle view. Here, changes in expectations, due to news (new information), is put forward as the primary driver of economic fluctuations, and linked to economic fundamentals, i.e., total factor productivity (Barsky and Sims (2012), and Blanchard et al. (2013)). An alternative view of narratives and their role in explaining economic fluctuations builds on the classical work of Pigou (1927) and Keynes (1936) on capturing the market’s animal spirits where changes in agents’ expectation can be totally self-fulfilling or not rooted in economic fundamentals at all. Such mechanisms have for example been highlighted by Shiller (2000), and recent work by Angeletos and La’O (2013).

Since changes in expectations are not directly observable, and since economic feedback loops easily can confound the cause and effect relationship, it is intrinsically difficult to discriminate between these two opposing views. Empirical investigations have therefore resorted to using various high frequency and hard to predict economic variables, e.g., asset prices or consumer sentiment (Beaudry and Portier (2006), Barsky and Sims (2012)), to approximate news and changes in expectations. In contrast, our approach permits the usage of a primary source of (potential) new information directly, namely the news narratives.

To this end, we build on the results presented in the previous section and partition the high dimensional news topic dataset into what we loosely call “propagators” and “initiators”. The “propagators” are news topics with a high centrality score in the graphical Granger causality network. Such narratives predict many of the other series, but are also themselves predicted by a large share of other news topics. In contrast, the “initiators” are more exogenous. At the extreme they are not predicted by any of the other series, but they do still themselves have predictive power for other narratives (confer Table 6). Thus, any unexpected changes in these less central parts of the network should be less likely to be due to potential feedback loops, and more likely to represent new information.

Building on this simple logic, and focusing on the US, Figure 7a plots the first principal component estimate of the five most “exogenous” US-based news topic time series, i.e.,

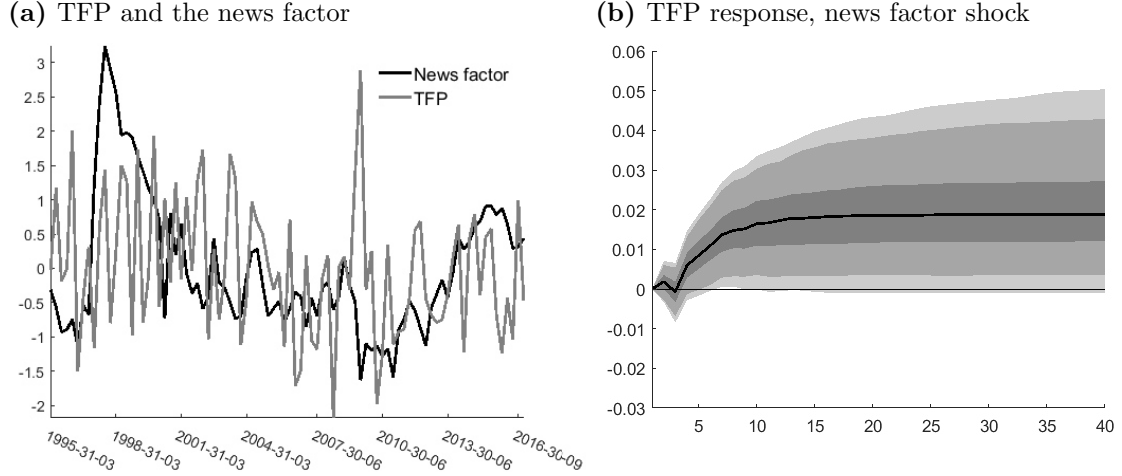


Figure 7. Figure 7a reports the estimated news factor together with TFP for the US. Figure 7b reports the response (in levels) of US TFP following a one standard deviation innovation in the news factor. The black solid line is the median estimate. The uncertainty bands reflect the 95, 90, and 50 percent quantiles, constructed from a residual bootstrap.

those with in degree equaling 0 from Table 6, together with total factor productivity (TFP). The factor estimate explains 55 percent of the total variation across the five variables, and is reported on a quarterly frequency. The TFP measure is adjusted for capacity utilization using the methodology suggested by Basu et al. (2006), and obtained from the *Federal Reserve Bank of San Francisco* web pages (Gerstein (2018)). As seen from the figure, the TFP estimate shows much more high frequency variation than the news factor. Still, there is a clear tendency for the two series to move together. Their contemporaneous correlation is 0.2.

To investigate the dynamic relationship between the news factor and TFP, and show how unexpected fluctuations in the news factor affects TFP, we formulate a simple bivariate Structural Vector Autoregression (SVAR) with these two variables. In the tradition of Beaudry and Portier (2006) and Barsky and Sims (2012), shocks to the news factor are identified using a recursive ordering where TFP is ordered first in the system and the news factor last. Thus, unexpected innovations in the news factor are orthogonal to contemporaneous TFP disturbances, and can only affect TFP with a lag. According to the news-driven business cycle view, and to the extent that shocks to the news factor contain new information, we expect a delayed but persistent increase in TFP. On the other hand, if the narratives just contain sentiment and noise, TFP should not respond at all to unexpected shocks in the news factor.

Figure 7b reports the cumulative response, i.e., the level, of TFP following a shock to the news factor. During the first year following the initial impulse, TFP is more or less unaffected. Then it increase significantly, and remains at a higher level than prior to the shock. This response pattern is as predicted by the news-driven business cycle view, and

Table 7. News factor and story examples. The story examples are found by querying the corpus for news articles where the five “initiator” news topics combined receive a high weight. Only the first sentences of each story are included in the table. The date of publication is printed in parenthesis.

(1998-04-07) Citrix Systems Inc. (CTXS) and Kronos Inc. (KRON) entered a joint agreement to market Citrix’s WinFrame software with Kronos’ Timekeeper C/S Version 2A. In a press release, Citrix said under the agreement Kronos has joined the Citrix Business Alliance, a coalition of vendors developing complementary products for its WinFrame thin-client/server software. Kronos provides systems that manage labor resources. Citrix Systems provides system software for thin-client/server computing...

(1998-04-14) The U.S. Commodity Futures Trading Commission Tuesday announced it will allow the Chicago Mercantile Exchange to trade the futures and options contracts aimed at providing risk management tools to credit card companies, banks and other consumer lending institutions...The CME quarterly Bankruptcy Index will be the world’s first futures and options to address default risk in the \$1.2 trillion consumer credit market, the CME said. The risk management tool could ultimately help lower consumer interest rates, the CME said...

(2014-06-11) Microsoft Corp.’s strategy for moving customers to its cloud email and productivity software is resonating with many corporate customers. Microsoft says the number of commercial seats for Office 365, its flagship productivity and email cloud service, more than doubled over the 12 months ending March 2014. It hopes its moves will lead to sales of a broader array of services to existing customers, including more complex business applications and cloud infrastructure services...

suggests that the news factor carries fundamental information, and not only noise and sentiment. The news shock also explains a large fraction of the variation in TFP. At the 10- and 40-quarter horizons, for example, as much as 22 and 52 percent of the variation in TFP can be attributed to the news shock.

Words clouds for the five narratives used to construct the news factor, the US-based topics *Clients*, *Software*, *Stocks*, *Bankruptcies*, and *Commodities*, are illustrated in Figure 16 in Appendix A. Examples of stories representative for these topics are reported in Table 7. As before, the narrative realism of the news topic-based approach stand out. The stories are clearly about technological changes, but also partly associated with developments in financial markets. However, as seen from Figure 17a, in Appendix A, the news factor does not work as a stand-in for surprising movements in asset markets. In particular, when we augment the SVAR model with quarterly returns from the Dow Jones Industrial Average, and order this variable above the news factor (but below TFP) in the recursively identified SVAR, our results remain basically unchanged from the benchmark case in Figure 7b.

The flip side of the argument used above is that unexpected innovations to the narratives with a high centrality score, i.e., the “propagators”, should be less likely to lead to a significant TFP response. Figure 17b confirms this hypothesis. When computing the first principal component of the two US-based news topic variables with the highest

centrality score, confer Table 6, and re-estimating the bivariate SVAR described above with this factor instead of the earlier news factor, we obtain insignificant results.

To the best of our knowledge, quarterly TFP statistics do not exist for Japan and the euro area (and, due to data availability they are hard to construct). Still, using interpolated quarterly TFP estimates based on the yearly statistics provided by the *European Commission*, we can get an impression of whether or not shocks to the US-based news factor tend to affect productivity levels globally as well. The results from this experiment are reported in Figure 18 in Appendix A. Following a news shock, the level of TFP in the euro area increases significantly, in line with the results for the US, although with a substantial lag of up to two years. For Japan, however, we get insignificant results. In that respect, it is interesting to note that among the 88 outgoing edges from the five US-based initiators used to construct the news factor (confer Table 6), 60 percent go directly to European news topics. Thus, in line with earlier results, there seems to be a stronger relationship between the US and Europe, than with the US and Japan, also when it comes to narratives associated with economic fundamentals.¹⁵

While our results clearly suggest that narratives, or at least some of them, carry fundamental information, we can not rightfully argue that these narratives cause TFP. There are well known potential problems with using SVAR models to try to uncover the structural effects of anticipated shocks (news shocks) (Sims and Zha (2006), Forni et al. (2017), Blanchard et al. (2013)). More broadly, establishing a causal relationship between narratives and economic developments, in terms of potential outcomes (Rubin (2005)), is difficult because of the obvious simultaneity between economic events and media coverage of the same events. Without some truly exogenous information, decoupling the effect of the new information component (the economic event) from the effect of the ether (the media generating the narrative or reporting on the event) is challenging.

Still, our results are very much in line with other newer studies trying to understand the underlying relationship between news and economic fluctuations using exogenous events and high-frequency data. For example, Larsen and Thorsrud (2017) use an exogenous strike in the newspaper market to show that up to 40 percent of the predictive effect from news topics to daily asset returns can be attributed to the causal effect of the media itself. Similarly, it is interesting to note that the narratives defined as “initiators” here overlap well in theme and meaning with the news topics associated with productivity

¹⁵At the 40-quarter horizons, 47 and 7 percent of the variation in the euro area and Japanese TFP measures, respectively, can be attributed to the news shock. The close to idiosyncratic behavior of Japanese productivity growth is also found in Crucini et al. (2011). They compute a common (yearly) component of productivity growth across G7 countries, and document that as little as 16 percent of the variation in TFP in Japan can be attributed to a common global component. In contrast, for the US this number is 43 percent.

for the Norwegian economy in [Larsen and Thorsrud \(2018\)](#). In that study, using a very different approach, news topics labeled *Funding*, *Stock market*, and *IT/startup* are among the most influential. These narratives share many important words with in particular the *Bankruptcies*, *Stocks*, and *Software* topics found to be important here.

5 Conclusion

To what extent are narratives informative for describing business cycle variation, do they go viral, how do they interact with each other, and are they associated with economic fundamentals or better understood as capturing the market’s animal spirits?

In this article we focus on the three major economies the US, Japan, and the euro area, and show how unstructured textual news data can be used to provide quantitative answers to these questions. We do so by first constructing daily business cycle indexes computed on the basis of the news topics the media writes about. We then derive virality indexes capturing the extent to which narratives relevant for growth go viral and affect economic fluctuations across borders, and finally use so called “Graphical Granger causality” modeling to cast light on cross-country narrative spillovers and whether or not narratives carry news or noise.

The resulting coincident indexes are shown to classify the phases of the cycle with high precision. At a broad level, the most important news narratives are shown to be associated with general macroeconomic developments, finance, and (geo-)politics. However, a vast set of narratives contribute to our index estimates across time, especially in times of expansion. In times of trouble, narratives associated with economic fluctuations become more sparse. Likewise, we show that narratives do go viral, with an average epidemic duration of 4-5 months, but mostly so in times of trouble. Finally, while narratives interact in complicated ways, we document that some news topics are clearly associated with economic fundamentals, as predicted by the news-driven business cycle view. Other narratives, on the other hand, show no such relationship, and are likely better explained by classical work capturing the market’s animal spirits.

More than providing definite answers, we offer a number of new results about the relationship between business cycles and narratives, and an analytical framework for quantifying such interactions. Natural extensions to the approach taken here include: comparing the topic model approach to other automated Natural Language Processing techniques and further investigate how textual data can be translated into useful time series; exploiting the high frequency nature of the news data and natural frictions in information flow (e.g., time zones), to design experiments better suited for understanding the underlying structural relationship between narratives and economic fluctuations.

References

- Altissimo, F., R. Cristadoro, M. Forni, M. Lippi, and G. Veronese (2010). New Eurocoin: Tracking Economic Growth in Real Time. *The Review of Economics and Statistics* 92(4), 1024–1034.
- Angeletos, G.-M. and J. La'O (2013). Sentiments. *Econometrica* 81(2), 739–779.
- Arnold, A., Y. Liu, and N. Abe (2007). Temporal causal modeling with graphical granger methods. In *Proceedings of the 13th ACM SIGKDD international conference on Knowledge discovery and data mining*, pp. 66–75. ACM.
- Aruoba, S. B., F. X. Diebold, and C. Scotti (2009). Real-Time Measurement of Business Conditions. *Journal of Business & Economic Statistics* 27(4), 417–427.
- Bai, J. and S. Ng (2013). Principal components estimation and identification of static factors. *Journal of Econometrics* 176(1), 18 – 29.
- Bai, J. and P. Wang (2014). Identification theory for high dimensional static and dynamic factor models. *Journal of Econometrics* 178(2), 794–804.
- Baker, S. R., N. Bloom, and S. J. Davis (2016). Measuring economic policy uncertainty. *The Quarterly Journal of Economics* 131(4), 1593–1636.
- Balke, N. S., M. Fulmer, and R. Zhang (2017). Incorporating the beige book into a quantitative index of economic activity. *Journal of Forecasting* 36(5), 497–514.
- Banbura, M., D. Giannone, M. Modugno, and L. Reichlin (2013). Now-casting and the real-time data flow. Working Paper Series 1564, European Central Bank.
- Barsky, R. B. and E. R. Sims (2012). Information, Animal Spirits, and the Meaning of Innovations in Consumer Confidence. *American Economic Review* 102(4), 1343–77.
- Basu, S., J. G. Fernald, and M. S. Kimball (2006). Are technology improvements contractionary? *American Economic Review* 96(5), 1418–1448.
- Beaudry, P. and F. Portier (2006). Stock Prices, News, and Economic Fluctuations. *American Economic Review* 96(4), 1293–1307.
- Bholat, D., S. Hansen, P. Santos, and C. Schonhardt-Bailey (2015). Text mining for central banks: Handbook. *Centre for Central Banking Studies* 33, pp. 1–19.
- Blanchard, O. J., J.-P. L’Huillier, and G. Lorenzoni (2013). News, Noise, and Fluctuations: An Empirical Exploration. *American Economic Review* 103(7), 3045–70.

- Blei, D. M. (2012). Probabilistic topic models. *Communications of the ACM* 55, 77–84.
- Blei, D. M., A. Y. Ng, and M. I. Jordan (2003). Latent Dirichlet Allocation. *J. Mach. Learn. Res.* 3, 993–1022.
- Bruner, J. (1991). The narrative construction of reality. *Critical inquiry* 18(1), 1–21.
- Burns, A. F. and W. C. Mitchell (1946). *Measuring Business Cycles*. Number 46-1 in NBER Books. National Bureau of Economic Research, Inc.
- Carter, C. K. and R. Kohn (1994). On Gibbs Sampling for State Space Models. *Biometrika* 81(3), 541–553.
- Chang, J., S. Gerrish, C. Wang, J. L. Boyd-graber, and D. M. Blei (2009). Reading tea leaves: How humans interpret topic models. In Y. Bengio, D. Schuurmans, J. Lafferty, C. Williams, and A. Culotta (Eds.), *Advances in Neural Information Processing Systems* 22, pp. 288–296. Curran Associates, Inc.
- Crucini, M., A. Kose, and C. Otrok (2011). What are the driving forces of international business cycles? *Review of Economic Dynamics* 14(1), 156–175.
- Dougal, C., J. Engelberg, D. Garcia, and C. A. Parsons (2012). Journalists and the stock market. *Review of Financial Studies* 25(3), 639–679.
- Forni, M., L. Gambetti, M. Lippi, and L. Sala (2017). Noisy news in business cycles. *American Economic Journal: Macroeconomics* 9(4), 122–52.
- Foroni, C. and M. Marcellino (2013). A survey of econometric methods for mixed-frequency data. Working Paper 2013/06, Norges Bank.
- Freeman, L. C. (1977). A set of measures of centrality based on betweenness. *Sociometry* 40(1), 35–41.
- Gentzkow, M., B. T. Kelly, and M. Taddy (2017). Text as data. Working Paper 23276, National Bureau of Economic Research.
- Gerstein, N. (2018). Total factor productivity. <https://www.frbsf.org/economic-research/indicators-data/total-factor-productivity-tfp/>. Accessed: 10.01.2017.
- Geweke, J. (1992). Evaluating the accuracy of sampling-based approaches to the calculation of posterior moments. In *In BAYESIAN STATISTICS*.

- Granger, C. W. J. (1969). Investigating causal relations by econometric models and cross-spectral methods. *Econometrica* 37(3), 424–38.
- Griffiths, T. L. and M. Steyvers (2004). Finding scientific topics. *Proceedings of the National academy of Sciences of the United States of America* 101(Suppl 1), 5228–5235.
- Hansen, S., M. McMahon, and A. Prat (2018). Transparency and Deliberation within the FOMC: A Computational Linguistics Approach. *The Quarterly Journal of Economics* (forthcoming).
- Harvey, A. C. (1990). *Forecasting, Structural Time Series Models and the Kalman Filter*. Cambridge Books. Cambridge University Press.
- Heinrich, G. (2009). Parameter estimation for text analysis. Technical report, Fraunhofer IGD.
- Kermack, W. O. and A. G. McKendrick (1927). A contribution to the mathematical theory of epidemics. *Proceedings of the Royal Society of London A: Mathematical, Physical and Engineering Sciences* 115(772), 700–721.
- Keynes, J. (1936). *The General Theory of Employment, Interest and Money*. London: MacMillan.
- Kim, S., N. Shephard, and S. Chib (1998). Stochastic Volatility: Likelihood Inference and Comparison with ARCH Models. *Review of Economic Studies* 65(3), 361–93.
- Koopman, S. J. and A. Harvey (2003). Computing observation weights for signal extraction and filtering. *Journal of Economic Dynamics and Control* 27(7), 1317 – 1333.
- Kose, A., C. Lakatos, F. L. Ohnsorge, and M. Stocker (2017). The global role of the U.S. economy: linkages, policies and spillovers. Policy Research Working Paper Series 7962, The World Bank.
- Kose, M. A., C. Otrok, and E. Prasad (2012). Global Business Cycles: Convergence Or Decoupling? *International Economic Review* 53(2), 511–538.
- Kose, M. A., C. Otrok, and C. H. Whiteman (2003). International business cycles: World, region, and country-specific factors. *American Economic Review* 93(4), 1216–1239.
- Larsen, V. H. and L. A. Thorsrud (2017). Asset returns, news topics, and media effects. Working Paper 2017/17, Norges Bank.

- Larsen, V. H. and L. A. Thorsrud (2018). The Value of News for Economic Developments. *Journal of Econometrics* (Forthcoming).
- Loughran, T. and B. McDonald (2011). When Is a Liability Not a Liability? Textual Analysis, Dictionaries, and 10-Ks. *The Journal of Finance* 66(1), 35–65.
- Lozano, A. C., N. Abe, Y. Liu, and S. Rosset (2009). Grouped graphical granger modeling for gene expression regulatory networks discovery. *Bioinformatics* 25(12), i110–i118.
- Marcellino, M., M. Porqueddu, and F. Venditti (2016). Short-Term GDP Forecasting With a Mixed-Frequency Dynamic Factor Model With Stochastic Volatility. *Journal of Business & Economic Statistics* 34(1), 118–127.
- Mariano, R. S. and Y. Murasawa (2003). A new coincident index of business cycles based on monthly and quarterly series. *Journal of Applied Econometrics* 18(4), 427–443.
- Mumtaz, H., S. Simonelli, and P. Surico (2011). International comovements, business cycle and inflation: A historical perspective. *Review of Economic Dynamics* 14(1), 176–198.
- Nakajima, J. and M. West (2013). Bayesian Analysis of Latent Threshold Dynamic Models. *Journal of Business & Economic Statistics* 31(2), 151–164.
- Nimark, K. P. and S. Pitschner (2016). Delegated information choice. CEPR Discussion Papers 11323, C.E.P.R. Discussion Papers.
- Nymand-Andersen, P. (2016). Big data: The hunt for timely insights and decision certainty. IFC Working Paper 14, Bank for International Settlements.
- Pang, B., L. Lee, and S. Vaithyanathan (2002). Thumbs up?: Sentiment classification using machine learning techniques. In *Proceedings of the ACL-02 Conference on Empirical Methods in Natural Language Processing - Volume 10*, EMNLP '02, Stroudsburg, PA, USA, pp. 79–86. Association for Computational Linguistics.
- Peress, J. (2014). The media and the diffusion of information in financial markets: Evidence from newspaper strikes. *The Journal of Finance* 69(5), 2007–2043.
- Phillips, P. C. B., S. Shi, and J. Yu (2015). Testing for multiple bubbles: Historical episodes of exuberance and collapse in the s&p 500. *International Economic Review* 56(4), 1043–1078.
- Pigou, A. (1927). *Industrial Fluctuations*. London: MacMillan.

- Ramey, V. (2016). *Macroeconomic Shocks and Their Propagation*, Volume 2 of *Handbook of Macroeconomics*, Chapter 0, pp. 71–162. Elsevier.
- Rubin, D. B. (2005). Causal inference using potential outcomes: Design, modeling, decisions. *Journal of the American Statistical Association* 100(469), 322–331.
- Scott, S. L. and H. R. Varian (2013). Bayesian Variable Selection for Nowcasting Economic Time Series. NBER Working Papers 19567, National Bureau of Economic Research, Inc.
- Shapiro, A. H., M. Sudhof, and D. J. Wilson (2017). Measuring News Sentiment. Working Paper Series 2017-1, Federal Reserve Bank of San Francisco.
- Shiller, R. J. (2000). *Irrational exuberance*. Princeton university press.
- Shiller, R. J. (2017). Narrative economics. *American Economic Review* 107(4), 967–1004.
- Shojaie, A. and G. Michailidis (2010). Discovering graphical granger causality using the truncating lasso penalty. *Bioinformatics* 26(18), i517–i523.
- Simon P. Anderson, J. W. and D. Strömberg (Eds.) (2015). Volume 1 of *Handbook of Media Economics*. North-Holland.
- Sims, C. A. (2003). Implications of rational inattention. *Journal of Monetary Economics* 50(3), 665 – 690. Swiss National Bank/Study Center Gerzensee Conference on Monetary Policy under Incomplete Information.
- Sims, C. A. and T. Zha (2006). Does Monetary Policy Generate Recessions? *Macroeconomic Dynamics* 10(02), 231–272.
- Stock, J. and M. Watson (2012). Disentangling the channels of the 2007-2009 recession. *Brookings Papers on Economic Activity Spring 2012*, 81–135.
- Stock, J. H. and M. W. Watson (1988). A Probability Model of The Coincident Economic Indicators. NBER Working Papers 2772, National Bureau of Economic Research, Inc.
- Stock, J. H. and M. W. Watson (2005). Understanding changes in international business cycle dynamics. *Journal of the European Economic Association* 3(5), 968–1006.
- Tetlock, P. C. (2007). Giving content to investor sentiment: The role of media in the stock market. *The Journal of Finance* 62(3), 1139–1168.
- Tetlock, P. C. (2015). Chapter 18 - the role of media in finance. In J. W. Simon P. Anderson and D. Strömberg (Eds.), *Handbook of Media Economics*, Volume 1 of *Handbook of Media Economics*, pp. 701 – 721. North-Holland.

- Tetlock, P. C., M. Saar-Tsechansky, and S. Macskassy (2008). More Than Words: Quantifying Language to Measure Firms' Fundamentals. *Journal of Finance* 63(3), 1437–1467.
- Thorsrud, L. A. (2016a). Nowcasting using news topics. Big Data versus big bank. Working Papers 46, Centre for Applied Macro- and Petroleum economics (CAMP), BI Norwegian Business School.
- Thorsrud, L. A. (2016b). Words are the new numbers: A newsy coincident index of business cycles. Working Papers 44, Centre for Applied Macro- and Petroleum economics (CAMP), BI Norwegian Business School.
- Travis, J. B. and s. Jordà (2011). Evaluating the Classification of Economic Activity into Recessions and Expansions. *American Economic Journal: Macroeconomics* 3(2), 246–77.
- Yuan, M. and Y. Lin (2006). Model selection and estimation in regression with grouped variables. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 68(1), 49–67.
- Zhao, L., H. Cui, X. Qiu, X. Wang, and J. Wang (2013). Sir rumor spreading model in the new media age. *Physica A: Statistical Mechanics and its Applications* 392(4), 995 – 1003.
- Zhou, X., J. Nakajima, and M. West (2014). Bayesian forecasting and portfolio decisions using dynamic dependent sparse factor models. *International Journal of Forecasting* 30, 963–980.

Appendices

Appendix A Additional results

Table 8. Granger causality tests and p-values. The news-based coincident indexes are aggregated to monthly series and included in a three variable Vector Autoregression (VAR). The estimation sample is 1996:M1 - 2016:M12, and we allow for three lags in the model (while our results are robust to both larger and smaller lag orders.).

<i>NCI-US</i>		<i>NCI-Japan</i>		<i>NCI-Euro</i>	
<i>NCI-Japan</i>	<i>NCI-Euro</i>	<i>NCI-US</i>	<i>NCI-Euro</i>	<i>NCI-US</i>	<i>NCI-Japan</i>
0.36	0.00	0.35	0.00	0.18	0.00

Table 9. US news topics. Subjective labeling and the most important words. Weights in parenthesis.

Id	Label	Top words (word probability)
0	Monetary policy	inflat, 0.056, monetari, 0.018, bernank, 0.018, technic, 0.015, greenspan, 0.014, resist, 0.012, minut, 0.011
1	Fiscal policy	budget, 0.059, save, 0.033, deficit, 0.028, balanc, 0.014, social, 0.014, ir, 0.013, trillion, 0.012, reduct
2	Education	school, 0.044, ms, 0.035, univers, 0.031, famili, 0.028, student, 0.024, educ, 0.022, children, 0.016, colleg
3	Funding	loan, 0.106, mortgag, 0.082, borrow, 0.029, lend, 0.025, articl, 0.024, analysi, 0.022, lender, 0.022, link
4	Entertainment	book, 0.014, film, 0.013, art, 0.013, music, 0.011, movi, 0.011, star, 0.01, play, 0.01, artist, 0.006, theater
5	Telecommunication	network, 0.048, wireless, 0.043, phone, 0.033, mobil, 0.029, telecom, 0.025, telecommun, 0.022, verizon, 0.022
6	Agriculture	edt, 0.062, corn, 0.014, est, 0.011, crop, 0.011, farmer, 0.01, ceo, 0.01, agricultur, 0.009, farm, 0.009
7	Environment	water, 0.031, fuel, 0.029, environment, 0.024, emiss, 0.016, clean, 0.016, solar, 0.014, renew, 0.014, wast
8	Strategy	strategi, 0.026, opportun, 0.019, expand, 0.018, success, 0.016, focu, 0.016, challeng, 0.013, focus, 0.012
9	Trading	vol, 0.361, avg, 0.167, ttl, 0.158, blk, 0.081, prev, 0.041, nm, 0.02, zero, 0.017, uptick, 0.017, nyse, 0.016
10	Pharmaceutical	drug, 0.097, pharmaceut, 0.024, treatment, 0.02, fda, 0.019, patient, 0.019, trial, 0.015, cancer, 0.014, studi
11	Media	media, 0.039, tv, 0.032, cabl, 0.029, advertis, 0.025, network, 0.025, warner, 0.023, televis, 0.023, broadcast
12	Petroleum	crude, 0.083, barrel, 0.073, gasolin, 0.041, inventori, 0.02, nymex, 0.019, gallon, 0.018, heat, 0.018, opec
13	Public safety	polic, 0.034, fire, 0.025, kill, 0.015, death, 0.013, protest, 0.013, man, 0.011, gun, 0.011, safeti, 0.01
14	Employment	employe, 0.08, worker, 0.076, union, 0.064, employ, 0.036, pension, 0.033, labor, 0.032, strike, 0.021, wage
15	Iraq	militari, 0.037, iraq, 0.036, war, 0.023, iraqi, 0.022, troop, 0.02, armi, 0.014, attack, 0.013, afghanistan
16	Market perform	merril, 0.012, usd, 0.011, neutral, 0.01, nasdaq, 0.01, tg, 0.01, ep, 0.01, valuat, 0.007, djia, 0.007
17	Health care	health, 0.116, care, 0.086, hospit, 0.028, medic, 0.026, insur, 0.021, medicar, 0.019, coverag, 0.016
18	News service	thomson, 0.039, guidanc, 0.036, exclud, 0.032, segment, 0.023, adjust, 0.023, item, 0.022, gross, 0.021, prior
19	Energy	electr, 0.069, util, 0.065, plant, 0.057, ga, 0.02, capac, 0.018, california, 0.013, facil, 0.012, transmiss
20	Natural gas	ga, 0.071, natur, 0.033, pipelin, 0.025, bp, 0.023, drill, 0.021, field, 0.019, explor, 0.016, refinari, 0.014
21	China	china, 0.152, chines, 0.058, asia, 0.037, hong, 0.025, kong, 0.023, asian, 0.022, beij, 0.019, export, 0.016
22	M&A	bid, 0.063, merger, 0.055, stake, 0.05, acquir, 0.032, transact, 0.025, combin, 0.021, takeov, 0.019, familiar
23	Advisory	trust, 0.053, brown, 0.031, bancorp, 0.018, advisor, 0.016, dj, 0.013, branch, 0.011, ohio, 0.011, mgmt, 0.01
24	Smartphones	appl, 0.053, devic, 0.025, iphon, 0.019, game, 0.018, phone, 0.016, mobil, 0.014, app, 0.013, smartphon, 0.012
25	Clients	client, 0.047, email, 0.043, assum, 0.036, dilut, 0.035, advis, 0.032, reader, 0.031, either, 0.029, along
26	Persuasion	know, 0.022, realli, 0.012, someth, 0.012, got, 0.011, happen, 0.011, cannot, 0.01, tell, 0.009, sure, 0.008
27	Elections	elect, 0.041, campaign, 0.033, obama, 0.031, democrat, 0.025, republican, 0.023, parti, 0.023, vote, 0.021
28	Software	softwar, 0.074, microsoft, 0.05, comput, 0.041, ibm, 0.02, network, 0.015, window, 0.015, oracl, 0.014, applic
29	Electronics	chip, 0.044, comput, 0.029, intel, 0.028, electron, 0.027, dell, 0.023, semiconductor, 0.021, equip, 0.018, pc
30	Regulations	regul, 0.061, sec, 0.044, practic, 0.018, regulatori, 0.016, law, 0.016, investig, 0.015, audit, 0.013, act
31	Food	food, 0.059, restaur, 0.028, brand, 0.021, chain, 0.014, mcdonald, 0.011, drink, 0.011, coffe, 0.01, cola, 0.01
32	Justice	court, 0.078, law, 0.036, judg, 0.029, lawsuit, 0.028, legal, 0.026, claim, 0.024, settlement, 0.023, appeal
33	Economic crisis	crisi, 0.028, reform, 0.019, imf, 0.018, institut, 0.017, emerg, 0.015, stabil, 0.014, commit, 0.011, rubin
34	Retail	brand, 0.03, mart, 0.025, wal, 0.024, chain, 0.022, holiday, 0.017, discount, 0.015, shop, 0.015, apparel, 0.014
35	Europe	london, 0.046, plc, 0.029, india, 0.024, franc, 0.022, french, 0.021, ag, 0.019, deutsch, 0.018, german, 0.017
36	Leadership	ceo, 0.034, vice, 0.033, serv, 0.025, join, 0.022, resign, 0.021, replac, 0.019, role, 0.019, appoint, 0.017
37	Terrorism	attack, 0.033, al, 0.023, terrorist, 0.021, terror, 0.019, israel, 0.017, palestinian, 0.014, pakistan, 0.013
38	Stocks	common, 0.086, symbol, 0.057, issuer, 0.054, regist, 0.041, titl, 0.035, filer, 0.031, ownership, 0.03, outstand
39	Health	test, 0.039, studi, 0.023, dr, 0.019, diseas, 0.016, human, 0.013, health, 0.012, patient, 0.011, cancer, 0.009
40	Insurance	insur, 0.121, life, 0.032, aig, 0.023, premium, 0.021, deposit, 0.018, re, 0.014, claim, 0.014, cover, 0.011
41	Russia	russia, 0.039, russian, 0.03, minist, 0.022, nato, 0.018, kosovo, 0.014, prime, 0.013, eu, 0.012, moscow, 0.01
42	Brokerage firms	morgan, 0.082, goldman, 0.056, stanley, 0.046, merril, 0.037, sach, 0.035, citigroup, 0.031, lynch, 0.03, ge
43	Stock indices	nasdaq, 0.063, composit, 0.03, nyse, 0.024, advanc, 0.023, lost, 0.023, poor, 0.016, climb, 0.016, ralli, 0.015
44	Documentation	letter, 0.048, review, 0.038, request, 0.025, document, 0.021, correct, 0.019, wrote, 0.017, sent, 0.017, respond
45	Internet	onlin, 0.042, googl, 0.04, internet, 0.03, search, 0.025, user, 0.025, yahoo, 0.023, facebook, 0.022, ventur
46	Commentary	blog, 0.044, david, 0.033, steven, 0.025, pm, 0.023, paul, 0.021, onlin, 0.021, miller, 0.015, morn, 0.012
47	The White House	bush, 0.069, white, 0.037, clinton, 0.034, georg, 0.02, secretari, 0.019, negoti, 0.016, leader, 0.014, free
48	East Asia	japan, 0.092, north, 0.077, south, 0.059, korea, 0.054, japanes, 0.046, dn, 0.033, tokyo, 0.029, korean, 0.024
49	Natural gas	ga, 0.04, natur, 0.035, weather, 0.03, la, 0.022, casino, 0.022, winter, 0.016, normal, 0.015, vega, 0.014
50	Currencies	euro, 0.085, currenc, 0.075, yen, 0.048, zone, 0.02, ecb, 0.013, japan, 0.011, london, 0.011, greec, 0.01, franc
51	Weapons	nuclear, 0.037, iraq, 0.035, iran, 0.033, weapon, 0.028, council, 0.028, sanction, 0.022, resolut, 0.016
52	Results	dec, 0.05, dividend, 0.048, aug, 0.044, sept, 0.043, asknewswir, 0.03, item, 0.029, exclud, 0.027, oct, 0.027
53	Volatility	auction, 0.036, hedg, 0.027, volatil, 0.022, spread, 0.017, fix, 0.017, dealer, 0.016, bid, 0.016, particip
54	Argumentation	often, 0.013, exampl, 0.012, rather, 0.021, approach, 0.01, actual, 0.009, fact, 0.008, argu, 0.008, studi, 0.007
55	Labor market	economist, 0.05, labor, 0.026, revis, 0.024, claim, 0.023, unemploy, 0.022, employ, 0.02, read, 0.019, payrol
56	Real estate	properti, 0.054, estat, 0.045, hotel, 0.031, squar, 0.02, center, 0.019, leas, 0.017, park, 0.015, owner, 0.013
57	Australia	mine, 0.033, australia, 0.025, chemic, 0.025, ltd, 0.024, materi, 0.021, coal, 0.02, australian, 0.019, ton
58	Fear	recess, 0.017, slow, 0.015, recoveri, 0.014, worri, 0.012, warn, 0.01, fear, 0.01, confid, 0.01, bad, 0.008
59	Events	yesterday, 0.019, bp, 0.013, yr, 0.013, ep, 0.012, jh, 0.011, djia, 0.011, pjv, 0.01, chanc, 0.01, dec, 0.008
60	California	san, 0.047, california, 0.046, counti, 0.031, calif, 0.027, lo, 0.027, angel, 0.027, francisco, 0.023, jersey
61	Bonds	moodi, 0.039, matur, 0.03, poor, 0.024, spread, 0.023, downgrad, 0.02, plu, 0.02, swap, 0.017, grade, 0.017
62	Market talk	edt, 0.038, kevin, 0.029, kingsburi, 0.025, est, 0.024, premarket, 0.021, ep, 0.012, ceo, 0.012, kevingingsburi
63	Latin America	mexico, 0.043, brazil, 0.036, de, 0.021, mexican, 0.02, latin, 0.019, brazilian, 0.017, local, 0.016, peso, 0.016
64	Automobiles	car, 0.061, auto, 0.056, gm, 0.048, vehicl, 0.047, ford, 0.038, motor, 0.036, truck, 0.019, chrysler, 0.017
65	Bankruptcies	bankruptci, 0.065, facil, 0.028, protect, 0.028, restructur, 0.027, creditor, 0.024, chapter, 0.019, court
66	Weather	storm, 0.023, hurrican, 0.022, florida, 0.019, west, 0.017, island, 0.016, damag, 0.015, coast, 0.015, texa
67	Sports	game, 0.035, team, 0.032, play, 0.02, season, 0.019, player, 0.018, sport, 0.014, win, 0.012, leagu, 0.01, coach
68	Investigations	investig, 0.037, alleg, 0.022, attorney, 0.02, prosecutor, 0.017, fraud, 0.014, crimin, 0.013, lawyer, 0.013
69	Aircrafts	defens, 0.035, boe, 0.034, aircraft, 0.027, engin, 0.024, air, 0.02, commerc, 0.017, jet, 0.016, space, 0.016
70	Options	option, 0.136, grant, 0.028, william, 0.028, lee, 0.021, tobacco, 0.02, exercis, 0.019, expir, 0.019, philip
71	Design	design, 0.012, ms, 0.009, room, 0.007, wear, 0.005, color, 0.005, style, 0.004, glass, 0.004, light, 0.004
72	Investing	portfolio, 0.038, mutual, 0.031, cap, 0.022, percent, 0.02, etf, 0.019, institut, 0.011, fidel, 0.011, track
73	Transportation	steel, 0.04, ship, 0.036, transport, 0.028, port, 0.018, train, 0.015, shipment, 0.014, rail, 0.012, pacif
74	Commodities	gold, 0.096, metal, 0.034, commod, 0.031, ounce, 0.026, silver, 0.025, copper, 0.022, chicago, 0.021, cme, 0.016
75	Aviation	airlin, 0.078, air, 0.037, flight, 0.035, travel, 0.034, carrier, 0.025, airport, 0.024, passeng, 0.02, pilot
76	Canada	canada, 0.059, canadian, 0.051, toronto, 0.016, td, 0.011, surpris, 0.009, jame, 0.008, ben, 0.008, whose, 0.008
77	Transactions	fee, 0.06, card, 0.055, access, 0.046, payment, 0.043, compens, 0.033, visit, 0.03, paid, 0.024, kit, 0.018
78	Congress	senat, 0.054, committe, 0.041, congress, 0.033, republican, 0.03, legisl, 0.03, vote, 0.03, democrat, 0.03
79	Medical equip.	johnson, 0.042, devic, 0.027, boston, 0.025, medic, 0.025, st, 0.017, heart, 0.016, scintif, 0.014, stent, 0.011

Table 10. Japan news topics. Subjective labeling and the most important words. Weights in parenthesis.

Id	Label	Top words (word probability)
0	Russia	russia, 0.07, russian, 0.043, crisi, 0.013, moscow, 0.011, whale, 0.01, rubl, 0.008, clinton, 0.008, seven, 0.008
1	Elections	koizumi, 0.048, elect, 0.045, polit, 0.03, vote, 0.029, poll, 0.015, junichiro, 0.015, reform, 0.015, win, 0.013
2	Wall Street	street, 0.091, wall, 0.086, journal, 0.06, stori, 0.036, wsj, 0.032, blog, 0.026, onlin, 0.016, real, 0.013
3	Aviation	airlin, 0.049, flight, 0.03, air, 0.026, ship, 0.02, travel, 0.018, transport, 0.017, carrier, 0.016, jal, 0.016
4	Persuasion	thing, 0.014, seem, 0.01, mean, 0.008, realli, 0.008, cannot, 0.008, might, 0.008, know, 0.008, tri, 0.007, someth
5	Industry	steel, 0.129, nippon, 0.055, chemic, 0.04, materi, 0.035, ton, 0.027, produc, 0.024, plant, 0.024, capac, 0.02
6	Fed	fed, 0.073, feder, 0.062, inflat, 0.053, hike, 0.02, eas, 0.014, greenspan, 0.014, committe, 0.013, chairman, 0.013
7	Internet	http, 0.071, www, 0.061, link, 0.049, web, 0.046, visit, 0.043, partner, 0.04, today, 0.038, access, 0.038, site
8	Insurance	insur, 0.141, life, 0.083, pension, 0.044, mutual, 0.027, return, 0.023, save, 0.019, marin, 0.018, premium, 0.018
9	Performance	impact, 0.075, small, 0.043, neg, 0.037, size, 0.031, margin, 0.029, affect, 0.028, perform, 0.024, factor, 0.023
10	US politics	administr, 0.038, secretari, 0.037, clinton, 0.034, washington, 0.034, treasury, 0.033, summer, 0.028, rubin, 0.024
11	Telecommunication	phone, 0.057, ntt, 0.056, mobil, 0.053, network, 0.036, commun, 0.033, telecommun, 0.033, telecom, 0.027, telephon
12	Fixed income	bill, 0.065, auction, 0.052, cash, 0.051, bid, 0.047, discount, 0.043, singapor, 0.027, deposit, 0.024, particip
13	Market performance	usd, 0.025, target, 0.023, hiroyuki, 0.019, resist, 0.019, break, 0.019, volatil, 0.018, technic, 0.018, kachi
14	M&A	sharehold, 0.052, stake, 0.049, bid, 0.032, acquisit, 0.031, equiti, 0.024, acquir, 0.021, valu, 0.02, board, 0.02
15	Employment	job, 0.082, worker, 0.049, labor, 0.044, employ, 0.042, employe, 0.028, union, 0.026, unemploy, 0.021, wage, 0.02
16	Australia	australia, 0.037, australian, 0.035, wsj, 0.029, zealand, 0.018, target, 0.017, au, 0.012, lion, 0.008, nz, 0.007
17	Economics data	survey, 0.068, economist, 0.057, surplu, 0.035, adjust, 0.034, gdp, 0.031, revis, 0.025, sentiment, 0.02, season
18	Justice	file, 0.044, court, 0.041, case, 0.031, claim, 0.02, protect, 0.019, rule, 0.019, settlement, 0.016, bankruptci
19	Months	dec, 0.061, bureau, 0.053, held, 0.051, oct, 0.049, sept, 0.042, jan, 0.041, nov, 0.039, feb, 0.035, aug, 0.034
20	Media	music, 0.03, soni, 0.027, movi, 0.019, broadcast, 0.017, film, 0.017, media, 0.015, tv, 0.014, televis, 0.013
21	Electronics	soni, 0.073, electron, 0.054, matsushita, 0.033, sharp, 0.027, digit, 0.023, tv, 0.023, display, 0.022, camera
22	Europe	europ, 0.089, germani, 0.051, franc, 0.048, german, 0.038, french, 0.033, itali, 0.021, london, 0.02, pari, 0.019
23	Pharmaceuticals	drug, 0.056, pharmaceut, 0.029, health, 0.018, medic, 0.018, approv, 0.017, patient, 0.016, treatment, 0.015, studi
24	Oil and gas	project, 0.077, ga, 0.075, oil, 0.035, natur, 0.034, energi, 0.026, field, 0.021, lng, 0.017, shell, 0.014
25	Market commentary	finish, 0.029, select, 0.028, osaka, 0.026, afternoon, 0.025, section, 0.021, unchang, 0.019, player, 0.018
26	Unknown	right, 0.056, name, 0.04, full, 0.028, home, 0.027, publish, 0.026, jame, 0.026, along, 0.024, send, 0.022, reader
27	Mining	mine, 0.046, ton, 0.032, gold, 0.031, metal, 0.029, coal, 0.027, copper, 0.026, iron, 0.024, produc, 0.024, ore
28	Outlook	recoveri, 0.085, outlook, 0.039, slow, 0.026, recov, 0.022, indic, 0.018, pace, 0.017, slowdown, 0.015, trend
29	America	america, 0.044, american, 0.039, brazil, 0.036, mexico, 0.03, emerg, 0.021, latin, 0.019, canada, 0.018, brazilian
30	Stimulus	packag, 0.077, reform, 0.058, stimulu, 0.048, hashimoto, 0.029, implement, 0.028, structur, 0.023, deregul, 0.019
31	South Asia	indonesia, 0.056, india, 0.049, singapor, 0.047, thailand, 0.043, malaysia, 0.029, southeast, 0.027, philippin
32	Economic crisis	crisi, 0.032, fear, 0.024, worri, 0.023, recess, 0.017, plung, 0.016, caus, 0.014, warn, 0.014, hurt, 0.013, emerg
33	Financial companies	nomura, 0.071, mean, 0.048, daiwa, 0.048, brokerag, 0.047, nikko, 0.041, morgan, 0.039, research, 0.038, stanley
34	Motor	motor, 0.098, toyota, 0.075, nissan, 0.062, car, 0.061, vehicl, 0.051, honda, 0.043, auto, 0.038, model, 0.02
35	Currencies	dealer, 0.092, quot, 0.046, player, 0.042, london, 0.036, sterl, 0.027, slightli, 0.021, deutsch, 0.019, est
36	Military	militari, 0.028, defens, 0.028, iraq, 0.027, attack, 0.026, forc, 0.024, war, 0.021, troop, 0.013, terrorist, 0.01
37	Fiscal policy	deficit, 0.033, balanc, 0.026, flow, 0.018, potenti, 0.013, gap, 0.013, shift, 0.012, signific, 0.012, reflect
38	Computer games	game, 0.09, nintendo, 0.025, soni, 0.022, consol, 0.02, microsoft, 0.016, playstat, 0.015, play, 0.015, xbox
39	Euro Zone	euro, 0.143, zone, 0.024, pound, 0.02, strategist, 0.017, franc, 0.013, swiss, 0.013, versu, 0.012, greenback
40	Negotiation	agreement, 0.066, negoti, 0.042, repres, 0.019, organ, 0.019, member, 0.018, tariff, 0.018, free, 0.016, wto
41	Korea	korea, 0.108, north, 0.103, south, 0.06, korean, 0.05, nuclear, 0.037, program, 0.017, seoul, 0.015, kim, 0.015
42	Retail	retail, 0.094, store, 0.085, softbank, 0.039, depart, 0.03, chain, 0.022, card, 0.02, custom, 0.02, sprint, 0.017
43	Agriculture	food, 0.039, beef, 0.023, beer, 0.018, case, 0.016, agricultur, 0.013, asahi, 0.013, ban, 0.013, kirin, 0.013
44	Mitsubishi	mitsubishi, 0.158, trust, 0.095, sumitomo, 0.087, mitsui, 0.058, merger, 0.033, ufj, 0.03, mizuho, 0.026, dai
45	Energy	power, 0.11, plant, 0.07, electr, 0.052, nuclear, 0.052, reactor, 0.026, energi, 0.024, util, 0.018, fuel, 0.017
46	Communication	spokesman, 0.097, ask, 0.052, discuss, 0.044, detail, 0.038, statement, 0.036, decid, 0.03, condit, 0.025, confirm
47	Stock listings	list, 0.071, section, 0.037, counter, 0.032, initi, 0.023, tse, 0.022, ipo, 0.022, finish, 0.018, limit, 0.018
48	Software	internet, 0.025, softwar, 0.023, appl, 0.022, comput, 0.018, user, 0.017, devic, 0.017, onlin, 0.015, yahoo, 0.012
49	Restructuring	restructur, 0.071, oversea, 0.062, divis, 0.04, subsidiari, 0.033, offic, 0.023, competit, 0.022, consolid, 0.022
50	Argumentation	might, 0.038, clear, 0.021, believ, 0.019, littl, 0.018, appear, 0.017, probabl, 0.017, though, 0.017, soon, 0.016
51	Competition	competit, 0.03, biggest, 0.016, strategi, 0.013, rival, 0.013, face, 0.011, expand, 0.011, small, 0.01, success
52	Bonds	yield, 0.094, treasury, 0.06, hedg, 0.02, equiti, 0.018, benchmark, 0.017, jgb, 0.016, portfolio, 0.014, fix
53	Market talk	wsj, 0.036, edt, 0.02, revenu, 0.016, cent, 0.012, today, 0.009, est, 0.008, kevin, 0.007, ceo, 0.007, premarket
54	Credit rating	basi, 0.048, moodi, 0.044, downgrad, 0.026, coupon, 0.026, matur, 0.023, standard, 0.021, denomin, 0.02, poor
55	Intervention	intervent, 0.055, vice, 0.028, miyazawa, 0.027, interven, 0.021, yuan, 0.021, sakakibara, 0.021, author, 0.02
56	Real estate	real, 0.078, construct, 0.063, build, 0.056, estat, 0.055, land, 0.039, project, 0.036, properti, 0.032, offic
57	Income	pretax, 0.117, parent, 0.086, revenu, 0.05, dividend, 0.037, bln, 0.035, consolid, 0.032, full, 0.032, revis
58	Natural disasters	earthquak, 0.023, area, 0.019, damag, 0.018, citi, 0.017, quak, 0.014, tsunami, 0.011, prefectur, 0.011, kilomet
59	Leadership	mr, 0.193, abe, 0.033, write, 0.017, offic, 0.014, former, 0.011, chairman, 0.011, shinzo, 0.01, appoint, 0.009
60	Petroleum	oil, 0.141, crude, 0.058, barrel, 0.036, energi, 0.021, refin, 0.02, light, 0.018, opec, 0.015, gasolin, 0.015
61	Automobiles	car, 0.061, auto, 0.05, vehicl, 0.039, gm, 0.038, ford, 0.035, motor, 0.031, toyota, 0.018, chrysler, 0.017, truck
62	Car technology	recal, 0.03, tire, 0.025, safeti, 0.019, batteri, 0.019, fuel, 0.019, vehicl, 0.018, hybrid, 0.015, test, 0.014
63	Transactions	purchas, 0.055, paper, 0.047, valu, 0.042, sold, 0.033, amount, 0.03, transact, 0.026, worth, 0.024, book, 0.023
64	Funding	loan, 0.158, bad, 0.051, lend, 0.038, borrow, 0.024, fail, 0.019, inject, 0.018, deposit, 0.017, lender, 0.017
65	Fiscal policy	tax, 0.17, spend, 0.093, budget, 0.078, incom, 0.059, consumpt, 0.028, revenu, 0.019, extra, 0.019, household
66	Alliances	joint, 0.111, ventur, 0.109, allianc, 0.047, stake, 0.044, tie, 0.039, partner, 0.033, form, 0.032, agreement
67	News	sourc, 0.131, kyodo, 0.076, cite, 0.064, newspap, 0.06, daili, 0.047, local, 0.045, decid, 0.032, saturday, 0.027
68	Monetary policy	boj, 0.117, eas, 0.05, board, 0.036, target, 0.027, governor, 0.024, deflat, 0.023, member, 0.021, takashi, 0.018
69	China	chines, 0.069, beij, 0.031, visit, 0.027, war, 0.024, island, 0.017, taiwan, 0.015, disput, 0.014, protest, 0.013
70	R&D	technolog, 0.034, research, 0.031, studi, 0.019, design, 0.018, univers, 0.015, creat, 0.014, center, 0.014, inform
71	Family	team, 0.008, famili, 0.008, women, 0.006, ms, 0.005, live, 0.005, young, 0.005, children, 0.004, home, 0.004, citi
72	Investigation	investig, 0.031, charg, 0.023, former, 0.02, scandal, 0.017, involv, 0.016, offic, 0.014, alleg, 0.014, arrest
73	IMF	imf, 0.051, crisi, 0.028, discuss, 0.017, emerg, 0.014, stabil, 0.013, aid, 0.013, cooper, 0.011, role, 0.01
74	Growth	jump, 0.039, oversea, 0.035, overal, 0.034, ago, 0.034, surg, 0.033, straight, 0.029, fourth, 0.028, climb, 0.023
75	Electro equipment	electr, 0.071, heav, 0.052, equip, 0.041, hitachi, 0.04, fuji, 0.036, machineri, 0.031, sanyo, 0.027, machin
76	Justice	propos, 0.038, requir, 0.032, regul, 0.031, allow, 0.027, law, 0.026, rule, 0.025, commiss, 0.024, approv, 0.023
77	Computer electronics	chip, 0.065, comput, 0.043, semiconductor, 0.04, nec, 0.036, technolog, 0.034, toshiba, 0.034, electron, 0.031
78	Hong Kong	hong, 0.056, kong, 0.056, australia, 0.016, shanghai, 0.016, composi, 0.014, reader, 0.013, name, 0.012, korea
79	Politics	parti, 0.12, democrat, 0.05, rule, 0.048, ldp, 0.047, liber, 0.046, opposit, 0.034, obuchi, 0.029, parliament

Table 11. Europe news topics. Subjective labeling and the most important words. Weights in parenthesis.

Id	Label	Top words (word probability)
0	Russia	russia, 0.109, russian, 0.093, moscow, 0.045, ukraine, 0.033, putin, 0.023, soviet, 0.014, rubl, 0.013, ukrainian
1	Automobiles	car, 0.073, vehicl, 0.043, auto, 0.04, motor, 0.029, gm, 0.024, ford, 0.019, volkswagen, 0.018, chrysler, 0.018
2	Oil drilling	bp, 0.075, poland, 0.031, polish, 0.019, spill, 0.016, safeti, 0.013, gulf, 0.012, rig, 0.012, warsaw, 0.012, zloti
3	Persuasion	might, 0.017, thing, 0.016, seem, 0.016, know, 0.012, question, 0.011, realli, 0.01, lot, 0.01, happen, 0.01, littl
4	Pharmaceuticals	drug, 0.068, patient, 0.022, pharmaceut, 0.021, treatment, 0.019, studi, 0.016, cancer, 0.014, trial, 0.011, fda
5	Macroeconomics	inflat, 0.069, economist, 0.042, survey, 0.033, manufactur, 0.018, statist, 0.018, revis, 0.014, zone, 0.013, export
6	Aviation	airlin, 0.076, air, 0.048, flight, 0.036, airport, 0.033, passeng, 0.029, carrier, 0.026, travel, 0.02, traffic
7	Schedule	dec, 0.049, jan, 0.044, nov, 0.038, date, 0.036, oct, 0.036, sourc, 0.033, ministri, 0.031, sept, 0.03, correct
8	Banks	pari, 0.105, de, 0.048, fr, 0.035, bnp, 0.021, pariba, 0.019, general, 0.017, societ, 0.016, jean, 0.014, suez
9	Investigation	investig, 0.049, charg, 0.033, alleg, 0.027, prosecutor, 0.018, former, 0.018, court, 0.013, probe, 0.012, trial
10	Regulations	regul, 0.052, review, 0.029, committe, 0.02, letter, 0.018, standard, 0.016, process, 0.014, regulatori, 0.013
11	HR	person, 0.019, relev, 0.015, scheme, 0.014, respect, 0.014, disclosur, 0.013, act, 0.013, date, 0.012, document
12	Aircrafts	aircraft, 0.03, defens, 0.03, engin, 0.03, boe, 0.03, airbu, 0.029, plane, 0.018, jet, 0.017, space, 0.015, ge
13	Commentary	partner, 0.074, journal, 0.062, access, 0.055, visit, 0.04, stori, 0.04, onlin, 0.026, blog, 0.024, isin, 0.023
14	Fiscal policy	budget, 0.071, deficit, 0.054, fiscal, 0.039, spend, 0.037, domest, 0.035, imf, 0.025, gross, 0.025, monetari, 0.025
15	Software	softwar, 0.028, nokia, 0.024, phone, 0.019, technolog, 0.018, devic, 0.017, microsoft, 0.017, comput, 0.017
16	Monetary policy	ecb, 0.065, zone, 0.061, grec, 0.059, greek, 0.036, bailout, 0.024, crisi, 0.021, sovereign, 0.016, monetari, 0.013
17	Real estate	citi, 0.044, home, 0.027, hotel, 0.026, properti, 0.024, train, 0.015, hous, 0.015, center, 0.015, real, 0.014, land
18	Nuclear	nuclear, 0.042, iran, 0.042, north, 0.035, korea, 0.027, council, 0.025, weapon, 0.023, sanction, 0.023, south
19	Shipping	south, 0.046, ship, 0.042, africa, 0.037, east, 0.034, north, 0.025, middl, 0.025, port, 0.021, african, 0.021
20	Nordic countries	paper, 0.036, free, 0.031, norway, 0.029, norwegian, 0.029, kroner, 0.027, danish, 0.025, denmark, 0.02, visit
21	Argumentation	small, 0.018, competit, 0.013, smaller, 0.011, attract, 0.01, research, 0.01, size, 0.009, exampl, 0.008, strategi
22	Sweden	swedish, 0.047, sweden, 0.039, ab, 0.039, stockholm, 0.037, kronor, 0.034, ericsson, 0.028, volvo, 0.014, man, 0.013
23	Margin	margin, 0.053, incom, 0.03, fourth, 0.026, divis, 0.025, charg, 0.019, exclud, 0.019, adjust, 0.019, item, 0.018
24	Germany	frankfurt, 0.053, berlin, 0.031, xe, 0.028, thoma, 0.019, merkel, 0.017, andrea, 0.014, commerzbank, 0.014
25	Energy	electr, 0.065, plant, 0.056, util, 0.04, water, 0.023, fuel, 0.018, emiss, 0.017, wind, 0.016, nuclear, 0.015, capac
26	Oil exploration	shell, 0.043, field, 0.043, ga, 0.036, explor, 0.032, reserv, 0.021, barrel, 0.017, block, 0.017, petroleum, 0.017
27	Retail	retail, 0.117, store, 0.063, chain, 0.021, brand, 0.02, shop, 0.014, custom, 0.013, supermarket, 0.012, luxuri, 0.01
28	Technology	ventur, 0.068, technolog, 0.064, joint, 0.062, manufactur, 0.035, electron, 0.027, siemen, 0.025, chip, 0.022, equip
29	Derivatives	option, 0.038, hedg, 0.036, deriv, 0.022, canada, 0.021, list, 0.021, canadian, 0.017, clear, 0.015, nasdaq, 0.014
30	Britain	pound, 0.082, british, 0.081, ireland, 0.041, britain, 0.038, irish, 0.035, sterl, 0.028, brown, 0.027, england
31	Crisis	crisi, 0.037, problem, 0.031, fear, 0.018, warn, 0.017, emerg, 0.014, worri, 0.013, collaps, 0.012, caus, 0.011
32	Latin America	brazil, 0.041, mexico, 0.026, de, 0.023, brazilian, 0.022, argentina, 0.021, america, 0.021, latin, 0.018, local
33	Mining	mine, 0.055, steel, 0.04, gold, 0.039, ton, 0.036, metal, 0.027, rio, 0.026, bhp, 0.025, copper, 0.023, miner, 0.023
34	Trading data	cent, 0.062, volum, 0.03, fiscal, 0.025, fourth, 0.021, thomson, 0.014, acquir, 0.012, guidanc, 0.009, jump, 0.008
35	Natural gas	ga, 0.13, project, 0.092, natur, 0.057, suppli, 0.039, pipelin, 0.038, export, 0.029, construct, 0.025, infrastruct
36	Persons	mr, 0.282, ms, 0.027, interview, 0.009, yesterday, 0.007, critic, 0.007, took, 0.007, person, 0.006, former, 0.006
37	M&A	bid, 0.092, merger, 0.056, takeov, 0.036, familiar, 0.025, person, 0.021, propos, 0.019, acquir, 0.018, combin
38	Outlook	impact, 0.028, reflect, 0.021, balanc, 0.019, neg, 0.018, factor, 0.018, view, 0.017, rel, 0.015, uncertainty, 0.013
39	Bonds	yield, 0.087, treasuri, 0.076, auction, 0.044, fix, 0.03, basi, 0.028, bid, 0.02, bill, 0.019, bundesbank, 0.018
40	Asia	china, 0.141, chines, 0.052, japan, 0.048, asia, 0.044, hong, 0.036, kong, 0.036, asian, 0.025, singapor, 0.021
41	On-line news	link, 0.096, front, 0.066, page, 0.06, analysi, 0.057, al, 0.046, commentari, 0.043, click, 0.038, rnd, 0.034
42	Education	famili, 0.015, school, 0.014, univers, 0.013, live, 0.009, student, 0.008, children, 0.008, educ, 0.007, women
43	Trading	trader, 0.088, dealer, 0.034, session, 0.034, volum, 0.032, quot, 0.025, dn, 0.023, vol, 0.021, vs, 0.018, morn
44	Switzerland	swiss, 0.111, ub, 0.053, zurich, 0.039, switzerland, 0.039, suiss, 0.024, martin, 0.016, abb, 0.015, client, 0.012
45	Market talk	edt, 0.027, est, 0.01, kevin, 0.009, ceo, 0.008, fed, 0.007, kingsburi, 0.007, amid, 0.006, yield, 0.006, ep, 0.006
46	Employment	job, 0.081, worker, 0.053, labor, 0.044, strike, 0.04, employ, 0.039, employe, 0.039, pension, 0.028, wage, 0.023
47	Spain	spain, 0.106, spanish, 0.074, madrid, 0.041, peseta, 0.024, de, 0.022, santand, 0.02, banco, 0.018, endesa, 0.016
48	Middle East	iraq, 0.053, turkey, 0.043, turkish, 0.023, iraqi, 0.019, war, 0.018, israel, 0.015, militari, 0.015, syria, 0.014
49	Media	media, 0.053, advertis, 0.032, tv, 0.027, televis, 0.026, broadcast, 0.022, music, 0.018, channel, 0.017, digit
50	Funding	loan, 0.093, mortgag, 0.041, lend, 0.034, lender, 0.029, liquid, 0.021, restructur, 0.019, balanc, 0.018, deposit
51	Emerging economies	emerg, 0.071, india, 0.06, austria, 0.025, local, 0.024, indian, 0.023, eastern, 0.021, austrian, 0.019, vienna
52	Credit rating	basi, 0.053, moodi, 0.029, fitch, 0.025, matur, 0.021, standard, 0.02, poor, 0.019, spread, 0.019, coupon, 0.017
53	Refineries	refineri, 0.035, refin, 0.021, tobacco, 0.021, shut, 0.019, texa, 0.016, mainten, 0.014, storm, 0.011, facil, 0.011
54	Australia	australia, 0.059, australian, 0.041, au, 0.025, zealand, 0.017, sydney, 0.013, andrew, 0.013, steward, 0.01, david
55	Italy	itali, 0.089, italian, 0.084, spa, 0.046, milan, 0.029, lire, 0.025, rome, 0.021, mi, 0.021, berlusconi, 0.015
56	Leadership	appoint, 0.032, ceo, 0.031, join, 0.024, former, 0.023, replac, 0.022, corpor, 0.021, role, 0.02, senior, 0.019
57	Transactions	exist, 0.019, custom, 0.019, opportun, 0.017, acquir, 0.016, approxim, 0.015, consider, 0.014, transact, 0.013
58	Investing	trend, 0.035, resist, 0.025, intraday, 0.023, technic, 0.022, reader, 0.022, chart, 0.022, bullish, 0.021, weekli
59	NATO	nato, 0.027, czech, 0.024, war, 0.023, republ, 0.019, kosovo, 0.016, serb, 0.015, allianc, 0.015, serbia, 0.014
60	Sports	mo, 0.032, game, 0.029, team, 0.024, play, 0.015, club, 0.015, sport, 0.014, win, 0.014, player, 0.014, match, 0.012
61	White House	hous, 0.034, administr, 0.03, washington, 0.028, american, 0.024, bush, 0.023, bill, 0.019, white, 0.016, obama
62	Brokerage firms	morgan, 0.058, barclay, 0.037, goldman, 0.031, stanley, 0.03, merrill, 0.027, ipo, 0.026, sach, 0.025, lynch, 0.024
63	Art	art, 0.009, design, 0.008, wine, 0.006, paint, 0.004, centuri, 0.004, collect, 0.004, artist, 0.004, gbp, 0.004
64	Petroleum	crude, 0.066, barrel, 0.06, cent, 0.032, brent, 0.022, gasolin, 0.022, suppli, 0.019, opec, 0.016, settl, 0.015
65	Public safety	polic, 0.02, fire, 0.018, rebel, 0.017, kill, 0.017, protest, 0.016, citi, 0.012, militari, 0.012, border, 0.01
66	Justice	court, 0.067, file, 0.047, claim, 0.03, legal, 0.026, law, 0.019, settlement, 0.019, appeal, 0.019, bankruptci
67	Taxation	tax, 0.145, payment, 0.046, dividend, 0.041, paid, 0.032, incom, 0.029, amount, 0.023, fee, 0.019, propos, 0.017
68	Food	food, 0.043, brand, 0.036, beer, 0.017, drink, 0.013, unilev, 0.01, nestl, 0.009, volum, 0.009, sugar, 0.009, brewer
69	Negotiation	negoti, 0.037, discuss, 0.031, leader, 0.026, confer, 0.02, summit, 0.02, organ, 0.015, cooper, 0.013, side, 0.011
70	Telecommunication	telecom, 0.061, mobil, 0.06, network, 0.056, telecommun, 0.033, wireless, 0.03, custom, 0.03, phone, 0.027, commun
71	Fear	weak, 0.033, outlook, 0.025, recoveri, 0.023, slow, 0.018, confid, 0.014, slowdown, 0.013, warn, 0.013, predict
72	Fed/BoJ	yen, 0.038, fed, 0.03, reserv, 0.02, japan, 0.016, strategist, 0.014, ralli, 0.013, japanes, 0.011, session, 0.009
73	Income	penc, 0.125, ln, 0.069, pretax, 0.022, dividend, 0.015, ftse, 0.014, upgrad, 0.013, neutral, 0.012, recommend, 0.012
74	EU	commiss, 0.101, eu, 0.077, brussel, 0.035, competit, 0.029, propos, 0.028, law, 0.02, commission, 0.013, regul
75	Insurance	insur, 0.117, life, 0.049, re, 0.028, premium, 0.027, lloyd, 0.023, standard, 0.02, rb, 0.019, royal, 0.016
76	Terrorism	attack, 0.047, terrorist, 0.021, terror, 0.021, polic, 0.02, suspect, 0.016, al, 0.016, bomb, 0.015, islam, 0.014
77	Benelux	dutch, 0.068, nv, 0.037, belgian, 0.035, netherland, 0.032, chemic, 0.031, amsterdam, 0.031, ing, 0.028, belgium
78	Elections	parti, 0.075, vote, 0.058, elect, 0.057, polit, 0.032, parliament, 0.026, prime, 0.022, opposit, 0.021, leader
79	Health	health, 0.033, test, 0.021, research, 0.018, agricultur, 0.012, vaccin, 0.012, human, 0.011, ban, 0.011, disease

Table 12. Best matching topics measured by the Jensen-Shannon divergence (JSD). The US based topics are used as the common “numeraire”.

US	Japan	Europe	US	Japan	Europe
Monetary policy	Fed	Fed/BoJ	Insurance	Insurance	Insurance
Fiscal policy	Fiscal policy	Fiscal policy	Russia	Russia	NATO
Education	Family	Education	Brokerage firms	Financial companies	Brokerage firms
Funding	Funding	Funding	Stock indices	Market commentary	Fed/BoJ
Entertainment	Family	Art	Documentation	Justice	HR
Telecommunication	Telecommunication	Telecommunication	Internet	Software	Software
Agriculture	Market talk	Market talk	Commentary	Persuasion	Persuasion
Environment	Energy	Energy	The White House	US politics	Negotiation
Strategy	Competition	Argumentation	East Asia	Korea	Nuclear
Trading	Market performance	Trading	Natural gas	Market talk	Petroleum
Pharmaceuticals	Pharmaceuticals	Pharmaceuticals	Currencies	Euro Zone	Fed/BoJ
Media	Media	Media	Weapons	Korea	Nuclear
Petroleum	Petroleum	Petroleum	Results	Months	Schedule
Public safety	Natural disasters	Public safety	Volatility	Bonds	Bonds
Employment	Employment	Employment	Argumentation	Persuasion	Persuasion
Iraq	Military	Middle East	Labor market	Economics data	Macroeconomics
Market performance	Market talk	Market talk	Real estate	Real estate	Real estate
Health care	Pharmaceuticals	Health	Australia	Mining	Mining
News service	Growth	Margin	Fear	Economic crisis	Crisis
Energy	Energy	Energy	Events	Market talk	Market talk
Natural gas	Oil and gas	Oil exploration	California	Justice	Real estate
China	South Asia	Asia	Bonds	Credit rating	Credit rating
M&A	M&A	M&A	Market talk	Market talk	Market talk
Advisory	Insurance	Brokerage firms	Latin America	America	Latin America
Smartphones	Software	Software	Automobiles	Automobiles	Automobiles
Clients	Unknown	Switzerland	Bankruptcies	Justice	Funding
Persuasion	Persuasion	Persuasion	Weather	Natural disasters	Refineries
Elections	Elections	Elections	Sports	Family	Sports
Software	Software	Software	Investigations	Investigation	Investigation
Electronics	Computer electronics	Technology	Aircrafts	Aviation	Aircrafts
Regulations	Justice	Regulations	Options	Transactions	Derivatives
Food	Agriculture	Food	Design	Family	Art
Justice	Justice	Justice	Investing	Bonds	Derivatives
Economic crisis	IMF	Fiscal policy	Transportation	Aviation	Shipping
Retail	Retail	Retail	Commodities	Mining	Fed/BoJ
Europe	Europe	Benelux	Aviation	Aviation	Aviation
Leadership	Leadership	Leadership	Canada	Fiscal policy	Outlook
Terrorism	Military	Terrorism	Transactions	Transactions	Taxation
Stocks	Transactions	HR	Congress	US politics	White House
Health	Pharmaceuticals	Health	Medical equip.	Pharmaceuticals	Pharmaceuticals

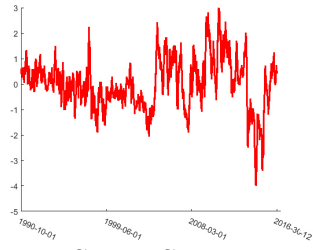
Table 13. Top five news topics across sub-samples for the *NCI-Japan* index. The *Example narratives* are found by querying the corpus for news articles where the five news topics listed in column two combined receive a high weight. Only the first sentences of each story are included in the table. The date of publication is printed in parenthesis.

	Top 5 news topics	Story example
1995 - 1999	Outlook Market Commentary Communication Europe Currencies	(1995-10-26) <i>The dollar is higher in early Tokyo trading Thursday than its levels late in New York Wednesday. Traders said that the yen's tone overall is weaker on rumors that Japanese investors may shift into mark-denominated investments when a large volume of Japanese government bonds mature Friday ...</i>
1999 - 2002	Outlook News Mitsubishi Fixed income Restructuring	(1999-02-13) <i>Toyo Trust & Banking Co. has agreed to transfer all its overseas securities custodian operations to Chase Manhattan (CMB) of the U.S., sources were quoted as saying in The Nihon Keizai Shimbun's Sunday edition. The accord represents Toyo Trust's complete withdrawal from overseas markets...</i>
2002 - 2006	Outlook Fixed income Investigation China Agriculture	(2002-05-12) <i>Japan plans to send a senior envoy to Beijing to negotiate the possible handover of five North Korean asylum seekers who were arrested by Chinese police last week on the grounds of a Japanese consulate in China, an official said Sunday... Video footage shot from a nearby building showed Chinese police rushing onto the grounds...</i>
2006 - 2009	Outlook Financial companies Russia Stock listings Negotiation	(2006-07-15) <i>After three days of nonstop negotiations, U.S. and Russian officials failed to seal a deal opening the way for Russia to join the World Trade Organization, dashing the Kremlin's hopes that the Group of Eight summit in St. Petersburg would showcase an agreement... Foreign banks, however, would still be barred from opening branches in Russia...</i>
2009 - 2013	Outlook Aviation Motor Unknown Natural disasters	(2010-04-20) <i>Nissan Motor Co. said Tuesday that the volcanic eruption in Iceland has forced it to temporarily suspend part of its domestic production lines as it is unable to airlift auto parts from Ireland... Nissan, which produced 2.74 million vehicles worldwide in 2009, expects Wednesday's stoppage to result in a production loss of 2,000 vehicles...</i>
2013 - 2016	Outlook Fed Market talk Car technology Wall Street	(2015-05-20) <i>Solid growth gives bank of Japan breathing room, though doubts linger after months of consistently undershooting expectations, Japan's economy actually outperformed forecasts in the first quarter... Chicago Fed President Charles Evans said Wednesday that it was by no means certain that the natural rate of unemployment in the U.S. is 5%...</i>

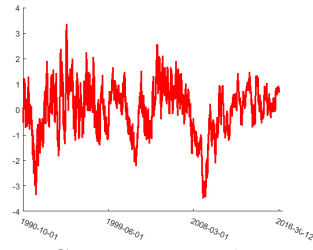
Table 14. Top five news topics across sub-samples for the *NCI-Euro* index. The *Example narratives* are found by querying the corpus for news articles where the five news topics listed in column two combined receive a high weight. Only the first sentences of each story are included in the table. The date of publication is printed in parenthesis.

	Top 5 news topics	Story example
1995 - 1999	Nordic countries	<i>(1995-12-11) Denmark's budget deficit is small and shrinking rapidly. Cutbacks in the welfare state hammered out two weeks ago between the government and the opposition haven't sparked mass protests... According to a recent Lehman Brothers survey of institutional investors ... a majority are overweighting Denmark,..</i>
	Switzerland	
	Fiscal policy	
	Emerging economics	
	Brokerage firms	
1999 - 2002	Australia	<i>(1999-01-22) Hoping to capitalize on U.S. investor interest in European buyouts, Morgan Grenfell Private Equity, Deutsche Bank AG's buyout unit, next week will begin marketing a EUR1.5 billion fund targeting acquisitions in Europe... particularly in Germany, which industry observers predict will be one of the main stages for the European M&A boom over the next few years...</i>
	Argumentation	
	Persons	
	M&A	
	Brokerage firms	
2002 - 2006	Middle East	<i>(2003-03-28) Crude oil futures relinquished early gains over the London morning Friday on a thin bout of profit taking, but gains are expected in later afternoon trade as people continue to price in a longer Iraq war than originally anticipated... U.S. Marines and Iraqi forces exchanged tank and artillery fire in Nasiriyah early Friday in a clash that set buildings in the city on fire...</i>
	Petroleum	
	Public safety	
	Trading	
	Shipping	
2006 - 2009	Macroeconomics	<i>(2008-06-16) ArcelorMittal (MT) is in a strong position to acquire Turkey's largest integrated steelmaker after increasing its stake in Turkish steel mill Erdemir to 24.98%, analysts said Monday... He said it makes more sense to increase a stake in a Turkish steel mill than build a new one from scratch since a new steel mill costs about \$1,500 a ton to build...</i>
	Middle East	
	Nuclear	
	Mining	
	M&A	
2009 - 2013	Macroeconomics	<i>(2011-10-20) A handful of companies sold debt Thursday, despite the continued distraction of European sovereign-debt worries. Three investment-grade issuers offered a combined \$2.5 billion in new debt while, in the junk-bond market, Kinetic Concepts Inc. (KCI) sold its \$2.3 billion term loan. Meanwhile, the municipal-bond market was fairly quiet Thursday,...</i>
	Trading data	
	Credit rating	
	Bonds	
	Funding	
2013 - 2016	Macroeconomics	<i>(2015-11-20) Eurozone consumers were more optimistic about their prospects in November, according to a survey by the European Commission that was largely completed before the Nov. 13 terror attacks on Paris... That possibility means the pickup in confidence is unlikely to dissuade policy makers at the European Central Bank from providing more stimulus when they meet in early December...</i>
	Sports	
	Asia	
	Monetary policy	
	Terrorism	

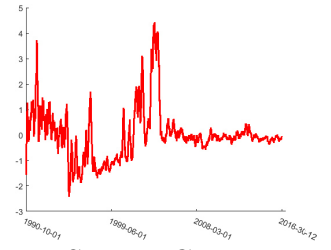
US T0: Monetary Policy



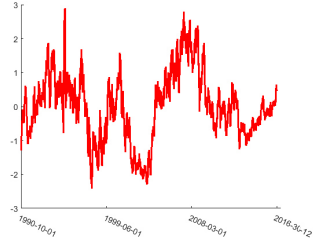
US T55: Labor market



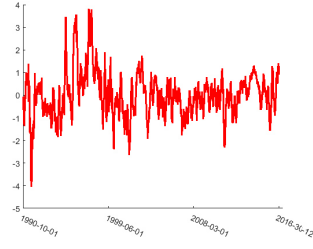
US T38: Stocks



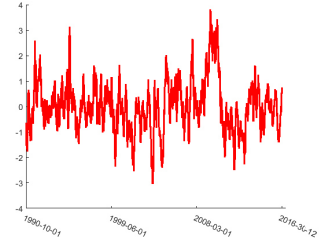
US T8: Strategy



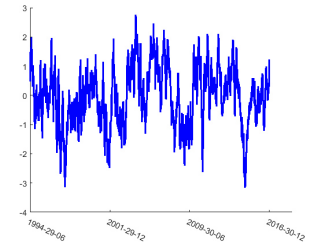
US T12: Petroleum



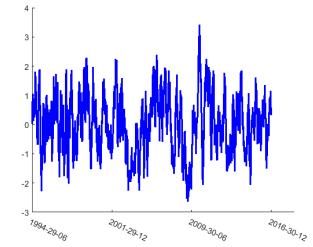
US T78: Congress



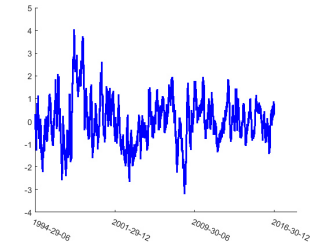
Japan T28: Outlook



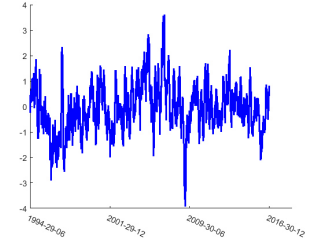
Japan T34: Motor



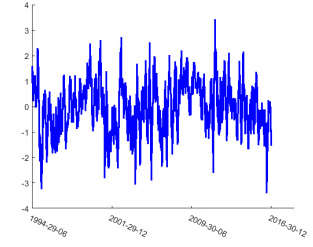
Japan T33: Financial companies



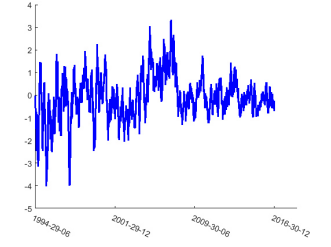
Japan T0: Russia



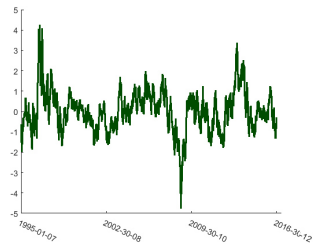
Japan T58: Natural disasters



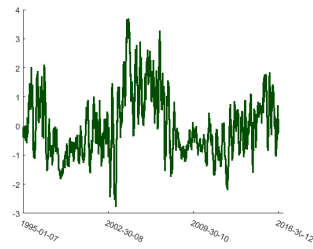
Japan T46: Communication



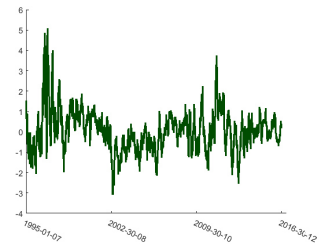
Europe T5: Macroeconomics



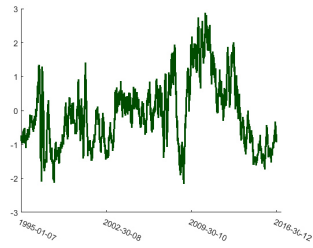
Europe T48: Middle East



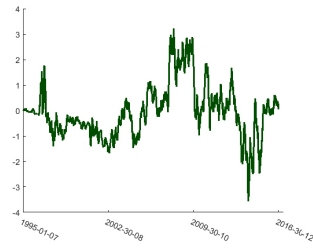
Europe T14: Fiscal policy



Europe T34: Trading data



Europe T58: Investing



Europe T79: Health

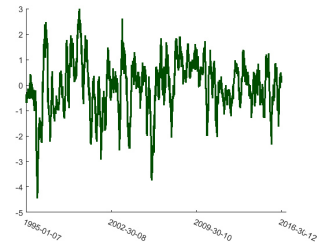
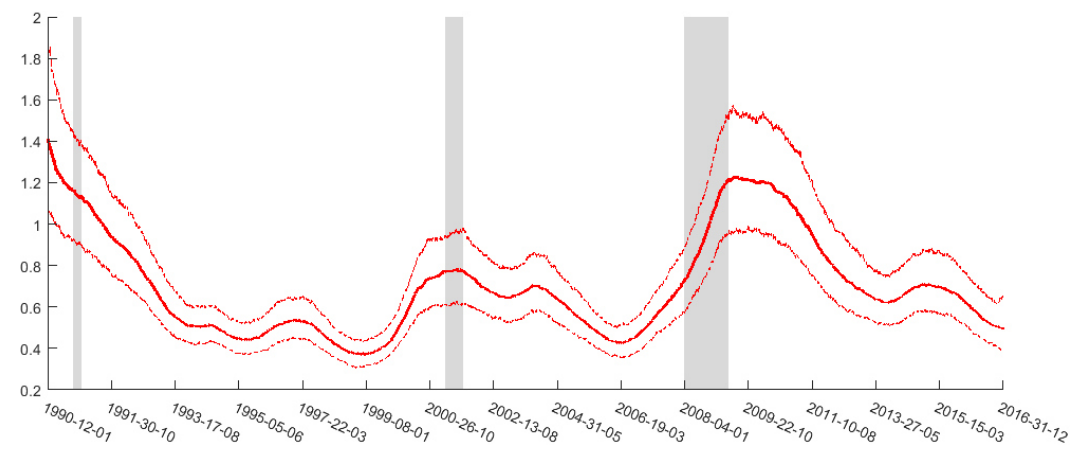
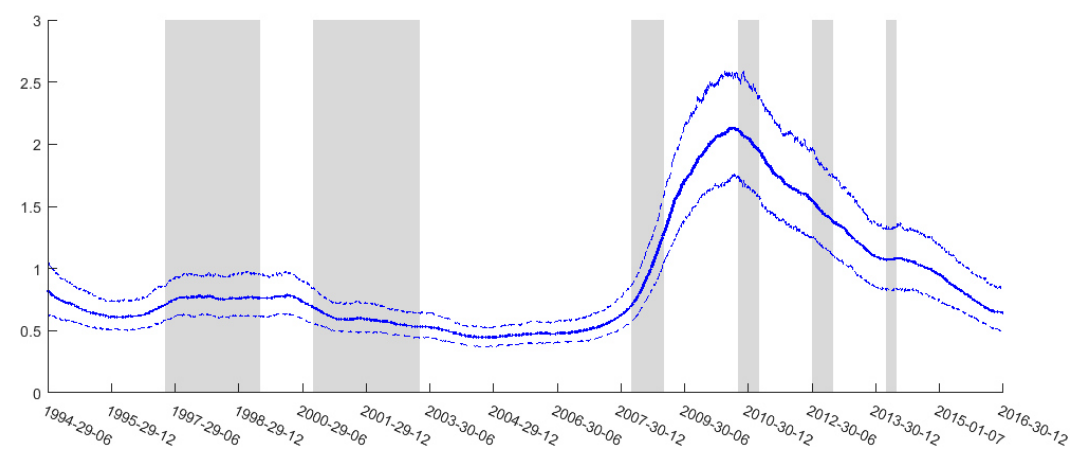


Figure 8. Daily news topic time series. See also Section 3.3, and Figure 2.

(a) SW-US



(b) SW-Japan



(c) SW-Europe

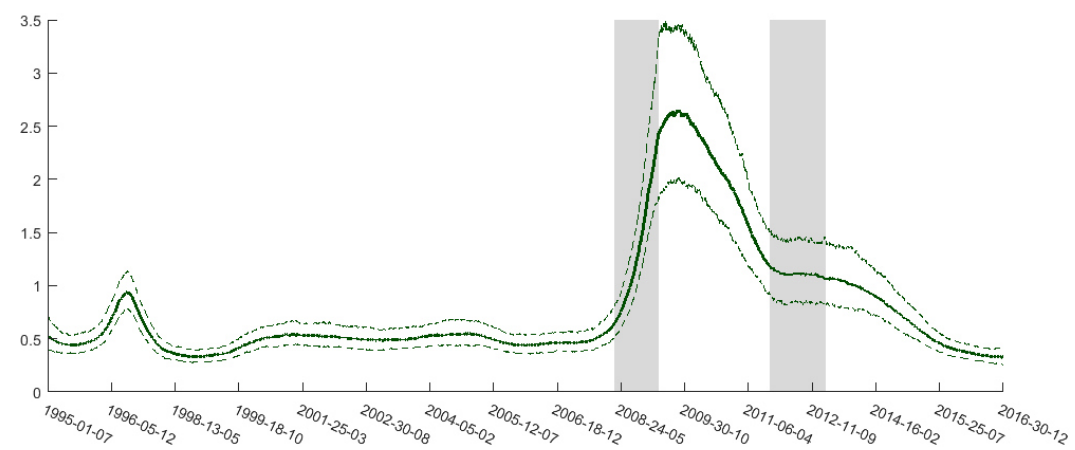


Figure 9. Standard error of the stochastic error in the daily coincident indexes. The colored solid line is the median, while the colored dotted lines are the 68 percent probability bands. The gray shaded areas illustrate recession periods as defined by NBER (US), ERCI (Japan), and CEPR (euro area).

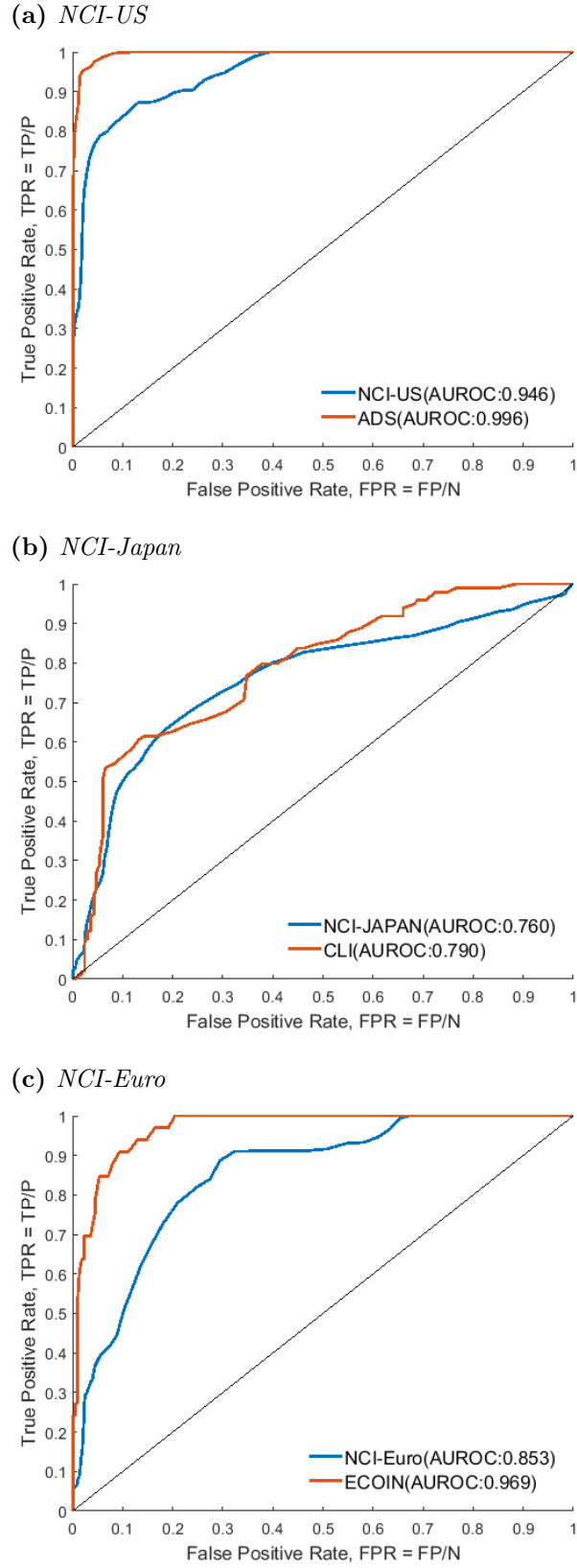


Figure 10. Receiver Operating Characteristics curves (ROC). As a measure of the unknown “truth” we use the business cycle phases defined by NBER (US), ERCI (Japan), and CEPR (euro area).

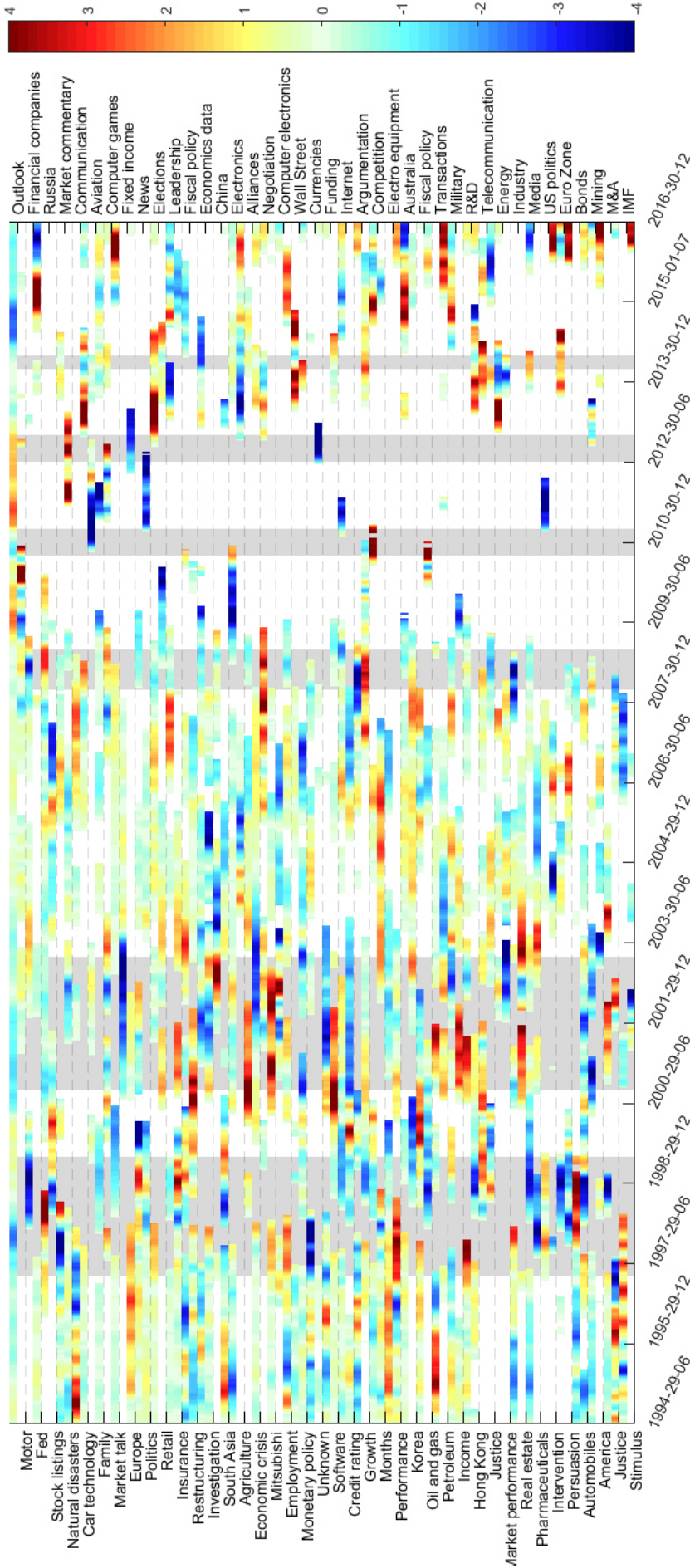


Figure 11. Japan: News topics and their contribution to *NCI* estimates across time. The reported decompositions are based on running the Kalman Filter using the posterior median estimates of the hyper-parameters and the time-varying factor loadings (at each time t). In the interest of readability, the topic names are reported on two y-axes with two-step increments. For example, the *Outlook* topic is associated with the first row (from above) in the figure, while the *Motor* topic is associated with the second row (from above). White areas illustrate the time-varying sparsity patterns. Recession periods, defined by ERCI, are illustrated using gray shading.

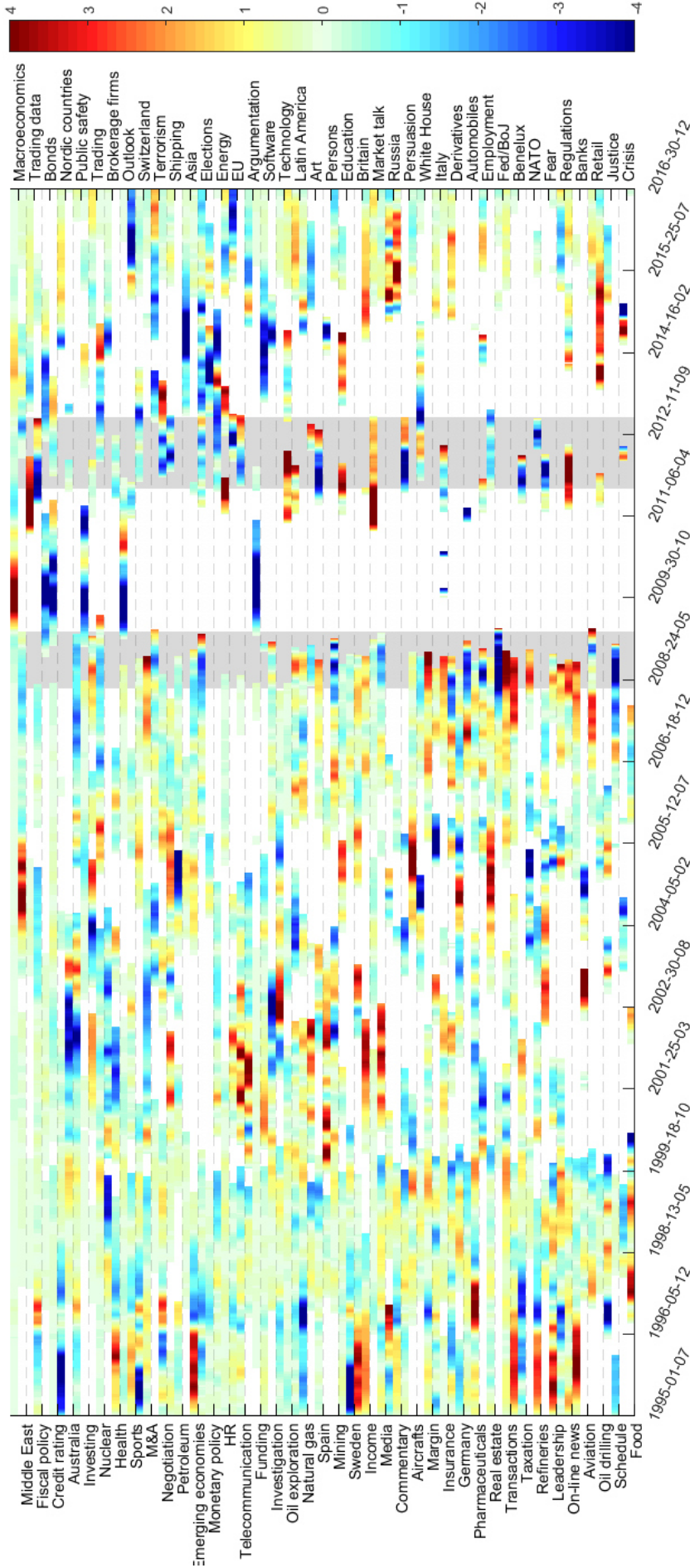


Figure 12. Europe: News topics and their contribution to *NCI* estimates across time. The reported decompositions are based on running the Kalman Filter using the posterior median estimates of the hyper-parameters and the time-varying factor loadings (at each time t). In the interest of readability, the topic names are reported on two y-axes with two-step increments. For example, the *Macroeconomics* topic is associated with the first row (from above) in the figure, while the *Middle East* topic is associated with the second row (from above). White areas illustrate the time-varying sparsity patterns. Recession periods, defined by CEPR, are illustrated using gray shading.

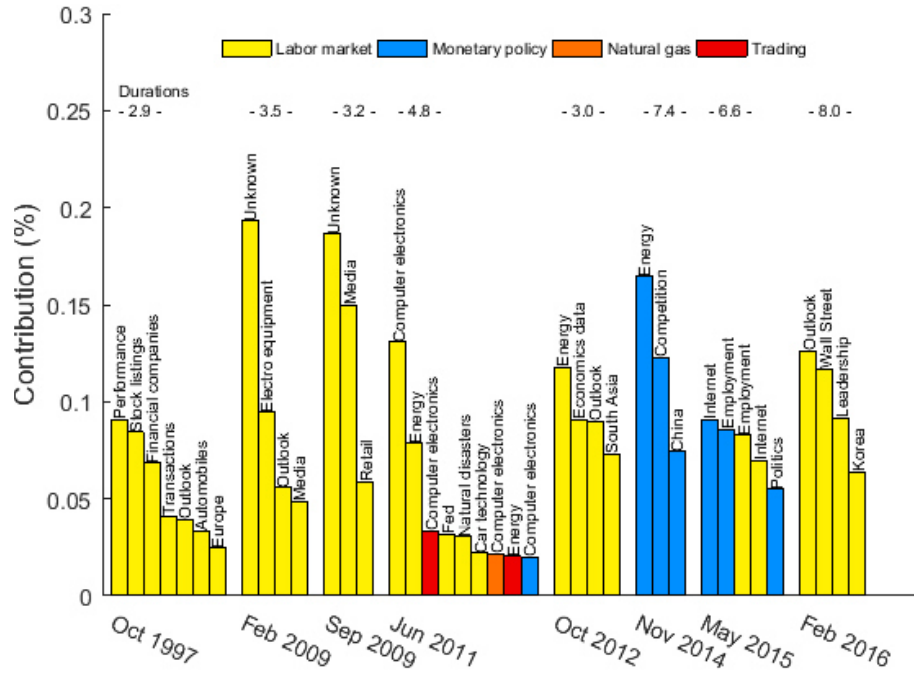
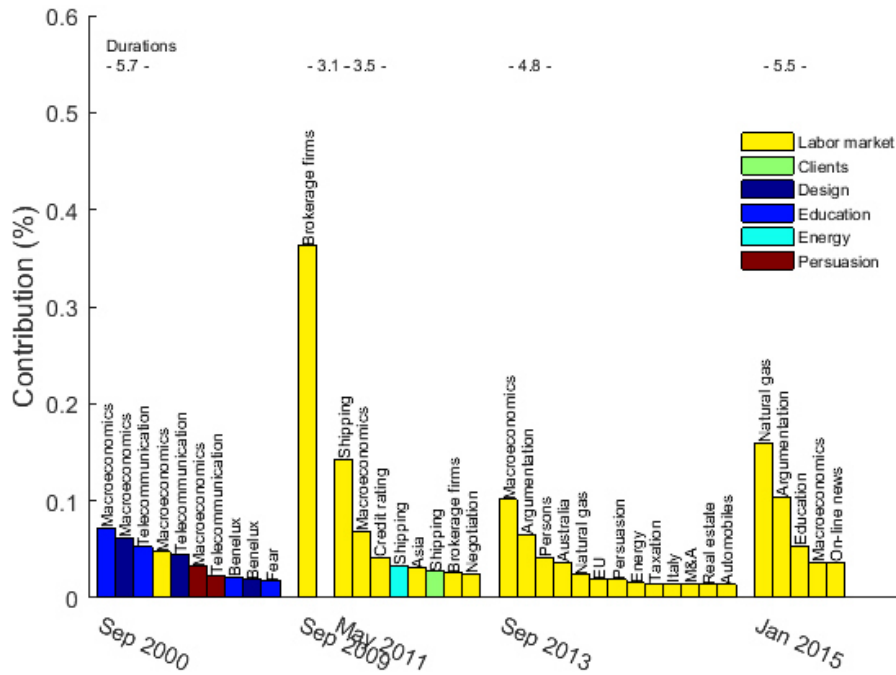
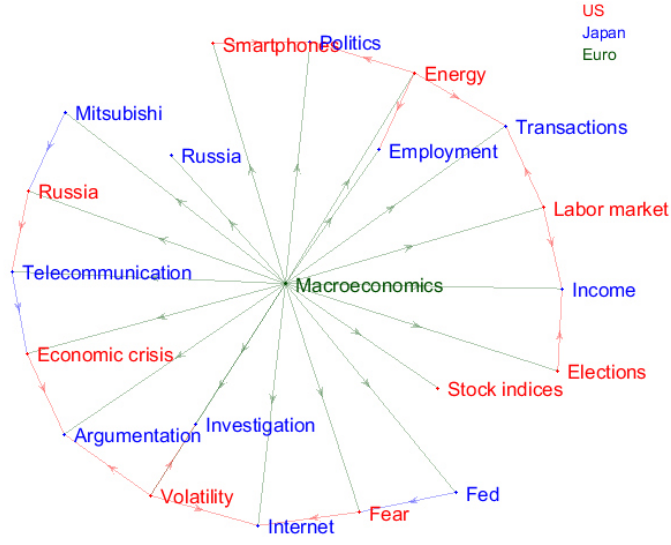
(a) $VIR^{Japan,US}$ (b) $VIR^{Euro,US}$ 

Figure 13. Epidemic periods and narratives. The peak dates and durations are calculated using a peak finder algorithm. Letting a 1 standard deviation increase (or more) in the indexes indicate that something goes viral, we define peaks as periods where the first derivative of the series equals 0. The duration of the epidemics are then estimated by a Gaussian distribution using the three coefficients from fitting a quadratic parabola to 7 data points centered at the peaks. For each epidemic period, we report the topic mappings that together explain up to 40 percent of the increases in the VIR indexes during the peak month. The legends, associated with the bar colors, report the name of the US-based topics, while the text above each bar report the associated Japan or Europe topic mapping.

(a) Europe T5 Macroeconomics



(b) Europe T48 Middle East

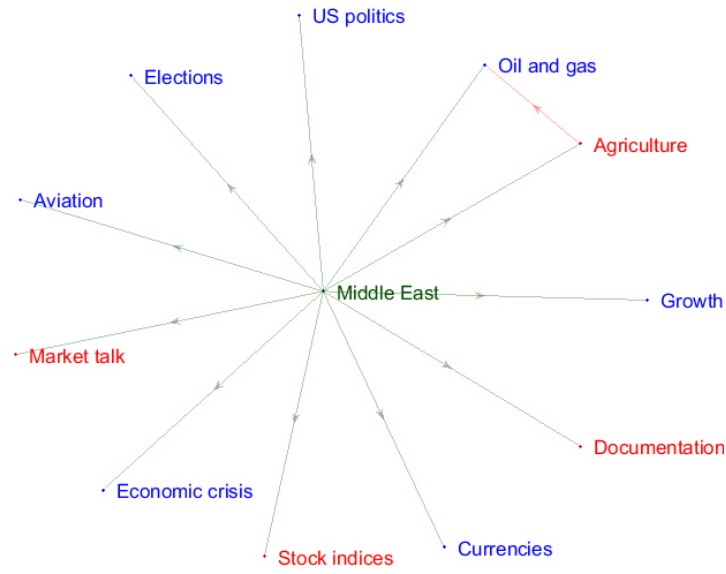
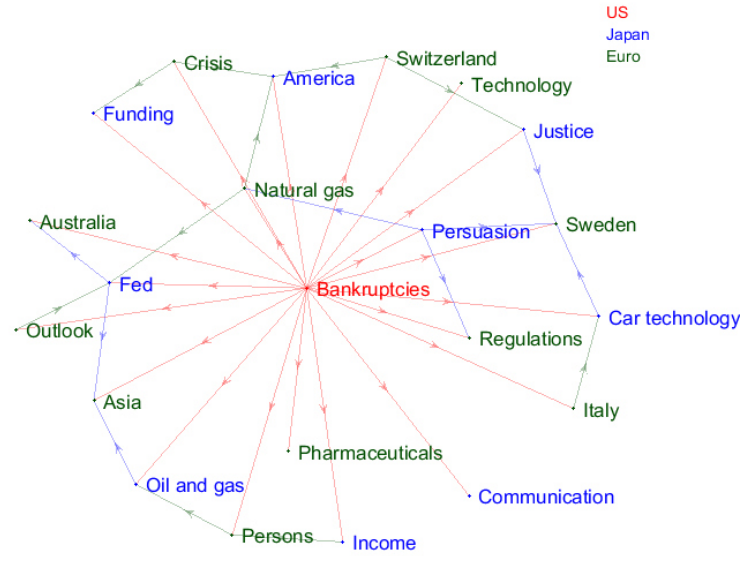


Figure 14. Network graph of the two most central narratives from the graphical Granger causality graph. The node and edge colors indicate from which country the topic belongs. In the interest of clarity, we only report the outgoing edges from the origin.

(a) US T65 Bankruptcies



(b) US T74 Commodities

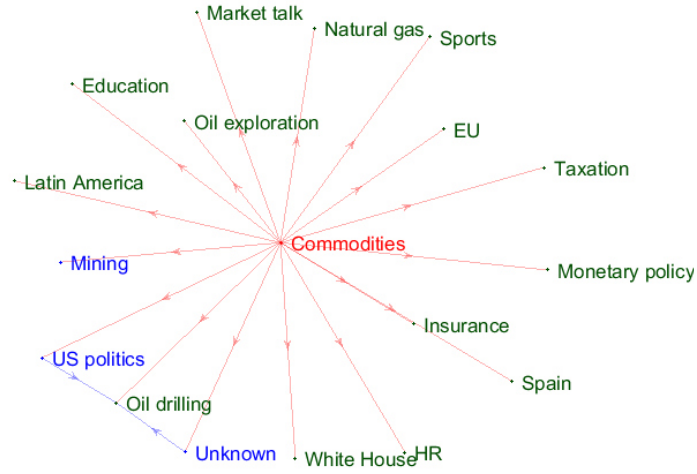
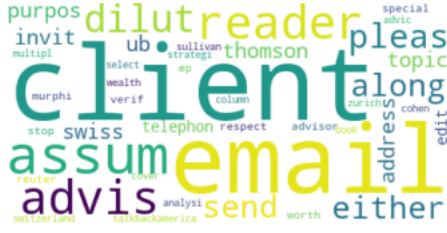


Figure 15. Network graph of the two least central narratives from the graphical Granger causality graph. The node and edge colors indicate from which country the topic belongs. In the interest of clarity, we only report the outgoing edges from the origin.

(a) US T25-Clients



(b) US T28-Software



(c) US T65-Bankruptcies

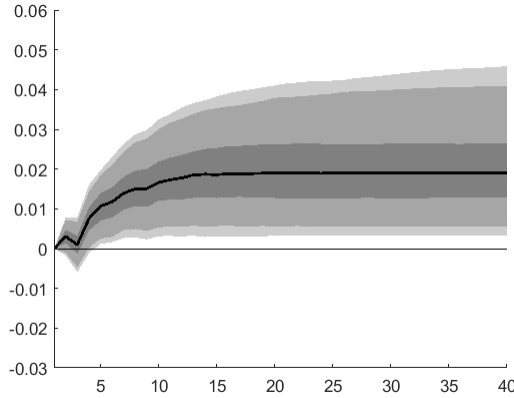


(d) US T74-Commodities



Figure 16. Word clouds and topic categorization of the “initiators” derived from Table 6 (The *Stocks* topic is reported in Figure 2). For each word cloud the size of a word reflects the probability of this word occurring in the topic. Each word cloud only contains a subset of all the words in the topic distribution.

(a) TFP response, with control



(b) TFP response, alternative factor

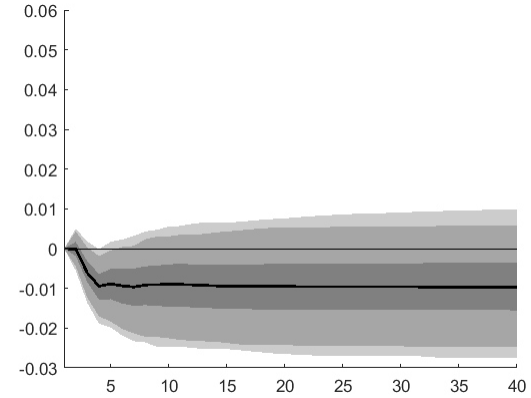
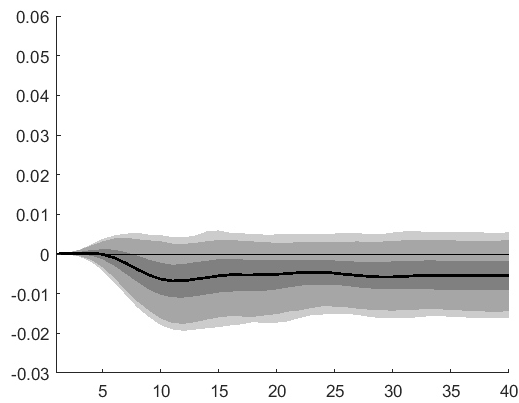


Figure 17. Figure 17a and 17b report the response (in levels) of US TFP following a one standard deviation innovation in a model controlling for asset returns, and when an alternative news factor is used, respectively. The black solid line is the median estimate. The uncertainty bands reflect the 95, 90, and 50 percent quantiles, constructed from a residual bootstrap.

(a) TFP response, Japan



(b) TFP response, euro area

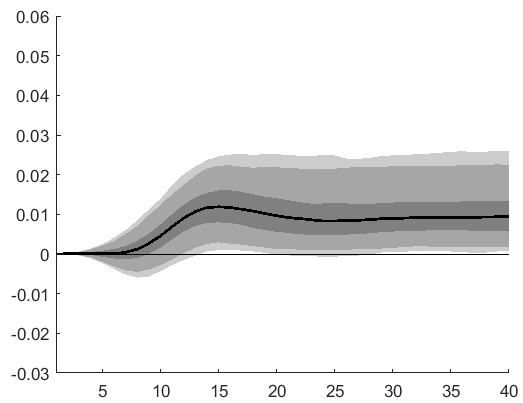


Figure 18. The figures report the response (in levels) of TFP following a one standard deviation innovation in the (US) news factor. In each impulse response graph, the black solid line is the median estimate. The uncertainty bands reflect the 95, 90, and 50 percent quantiles, constructed from a residual bootstrap.

Appendix B Reference classification

Because of the high dimensionality of the problem, and the fact that each of the estimated word distributions share words (although with different weight), it can be challenging to illustrate the output from the topic model. In addition, the corpus used for inference here is not publicly available, making it difficult for the reader to associate the estimated distributions with concrete examples. For this reason we investigate how the estimated topics relate to two external texts freely available to the public. These texts are the conclusion document from the US *Financial Crisis Inquiry Commision*, obtained from <https://fcic.law.stanford.edu/report/conclusions>, and the Federal Reserve Systems bi-annual *Monetary Policy Report to the Congress* from three different occasions, which can be downloaded from https://www.federalreserve.gov/monetarypolicy/mpr_default.htm.

These corpus are then first cleaned following the steps described in Section 3.1. Then, a procedure for querying documents outside the set on which the LDA is estimated is implemented, see Section C.2.

Table 15 summarizes the results. In short, when using the estimated topic distributions described in Section 3.2 to classify the conclusion document from the US *Financial Crisis Inquiry Commision* report, we find that the topics labeled *Funding*, *Economic crisis*, *Argumentation*, *Regulations*, and *Fear* together explains over 60 percent of the text. Thus, these topics are particularly associated with times of trouble, and also suggest that our subjective topic labeling is reasonable, although, perhaps, not perfectly descriptive.

Similarly, when classifying the Federal Reserve Systems bi-annual *Monetary Policy Report to the Congress*, we find that the topic labeled *Monetary policy* generally receives the highest probability (by far). However, across reports, and chairman, other topics also provide a good description. Examples are the *Labor market* and *Economic crisis* topics. Again, signaling that our subjective topic labeling is reasonable.

Table 15. Classification of alternative documents

Document	Date	Top news topics	Probability
The Financial Crisis Inquiry Report	January 27 2011	Funding	0.18
		Economic crisis	0.13
		Argumentation	0.12
		Regulations	0.10
		Fear	0.09
Testimony of Chairman Greenspan	July 18 1996	Monetary policy	0.18
		Labor market	0.16
		Argumentation	0.14
		Economic crisis	0.09
		Strategy	0.08
Testimony of Chairman Bernanke	July 21 2009	Economic crisis	0.26
		Monetary policy	0.12
		Regulations	0.10
		Fiscal policy	0.05
		Funding	0.05
Testimony of Chairman Yellen	June 21 2016	Monetary policy	0.36
		Labor market	0.19
		Economic crisis	0.07
		Fear	0.04
		Investing	0.04

Appendix C The textual data

Table 16. News article counts based on *Dow Jones* classification tags. Numbers are presented in percent of total articles in our sample. For example, 32 percent of the articles have a unique US tag, while 1 percent of the articles are tagged with the US and Japan identifier.

	US	Japan	Europe	US, Japan, Europe
US	0.32	0.01	0.03	
Japan		0.04	0.00	
Europe			0.08	
US, Japan, Europe				0.02

C.1 LDA estimation and specification

Figure 19 illustrates the LDA model graphically. The outer box, or plate, represent the whole corpus as M distinct documents (articles). $N = \sum_{m=1}^M N_m$ is the total number of words in all documents, and K is the total number of latent topics. Letting bold-

then to approximate the distribution:

$$P(\mathbf{Z}|\mathbf{W}; \alpha, \beta) = \frac{P(\mathbf{W}, \mathbf{Z}; \alpha, \beta)}{P(\mathbf{W}; \alpha, \beta)} \quad (14)$$

using Gibbs simulations, where α and β are the (hyper) parameters controlling the prior conjugate Dirichlet distributions for $\boldsymbol{\theta}_m$ and $\boldsymbol{\varphi}_k$, respectively. A very good explanation for how this method works is found in [Heinrich \(2009\)](#). The description below provides a brief summary only.

With the above definitions, the total probability of the model can be written as:

$$P(\mathbf{W}, \mathbf{Z}, \boldsymbol{\Theta}, \boldsymbol{\Phi}; \alpha, \beta) = \prod_{k=1}^K P(\boldsymbol{\varphi}_k; \beta) \prod_{m=1}^M P(\boldsymbol{\theta}_m; \alpha) \prod_{t=1}^N P(z_{m,t}|\boldsymbol{\theta}_m) P(w_{m,t}|\boldsymbol{\varphi}_{z_{m,t}}) \quad (15)$$

Integrating out the parameters $\boldsymbol{\varphi}$ and $\boldsymbol{\theta}$:

$$\begin{aligned} P(\mathbf{Z}, \mathbf{W}; \alpha, \beta) &= \int_{\boldsymbol{\Theta}} \int_{\boldsymbol{\Phi}} P(\mathbf{W}, \mathbf{Z}, \boldsymbol{\Theta}, \boldsymbol{\Phi}; \alpha, \beta) d\boldsymbol{\Phi} d\boldsymbol{\Theta} \\ &= \int_{\boldsymbol{\Phi}} \prod_{k=1}^K P(\boldsymbol{\varphi}_k; \beta) \prod_{m=1}^M \prod_{t=1}^N P(w_{m,t}|\boldsymbol{\varphi}_{z_{m,t}}) d\boldsymbol{\Phi} \int_{\boldsymbol{\Theta}} \prod_{m=1}^M P(\boldsymbol{\theta}_m; \alpha) \prod_{t=1}^N P(z_{m,t}|\boldsymbol{\theta}_m) d\boldsymbol{\Theta} \end{aligned} \quad (16)$$

In (16), the terms inside the first integral do not include a $\boldsymbol{\theta}$ term, and the terms inside the second integral do not include a $\boldsymbol{\varphi}$ term. Accordingly, the two terms can be solved separately. Exploiting the properties of the conjugate Dirichlet distribution it can be shown that:

$$\int_{\boldsymbol{\Theta}} \prod_{m=1}^M P(\boldsymbol{\theta}_m; \alpha) \prod_{t=1}^N P(z_{m,t}|\boldsymbol{\theta}_m) d\boldsymbol{\Theta} = \frac{\Gamma(\sum_{k=1}^K \alpha_k) \prod_{k=1}^K \Gamma(n_m^{(k)} + \alpha_k)}{\prod_{k=1}^K \Gamma(\alpha_k) \Gamma(\sum_{k=1}^K n_m^{(k)} + \alpha_k)} \quad (17)$$

and

$$\int_{\boldsymbol{\Phi}} \prod_{k=1}^K P(\boldsymbol{\varphi}_k; \beta) \prod_{m=1}^M \prod_{t=1}^N P(w_{m,t}|\boldsymbol{\varphi}_{z_{m,t}}) d\boldsymbol{\Phi} = \prod_{k=1}^K \frac{\Gamma(\sum_{t=1}^V \beta_t) \prod_{t=1}^V \Gamma(n_k^{(t)} + \beta_t)}{\prod_{t=1}^V \Gamma(\beta_t) \Gamma(\sum_{t=1}^V n_k^{(t)} + \beta_t)} \quad (18)$$

where $n_m^{(k)}$ denotes the number of word tokens in the m^{th} document assigned to the k^{th} topic, and $n_k^{(t)}$ is the number of times the t^{th} term in the vocabulary has been assigned to the k^{th} topic.

Since $P(\mathbf{W}; \alpha, \beta)$, in (14), is invariable for any of $\mathbf{Z}$, the conditional distribution $P(\mathbf{Z}|\mathbf{W}; \alpha, \beta)$ can be derived from $P(\mathbf{W}, \mathbf{Z}; \alpha, \beta)$ directly using Gibbs simulation and the conditional probability:

$$P(Z_{(m,n)} | \mathbf{Z}_{-(m,n)}, \mathbf{W}; \alpha, \beta) = \frac{P(Z_{(m,n)}, \mathbf{Z}_{-(m,n)}, \mathbf{W}; \alpha, \beta)}{P(\mathbf{Z}_{-(m,n)}, \mathbf{W}; \alpha, \beta)} \quad (19)$$

where $Z_{(m,n)}$ denotes the hidden variable of the n^{th} word token in the m^{th} document, and $\mathbf{Z}_{-(m,n)}$ denotes all $\mathbf{Z}$ s but $Z_{(m,n)}$. Denoting the index of a word token by $i =$

(m, n) , and using the expressions in (17) and (18), cancellation of terms (and some extra manipulations exploiting the properties of the gamma function) yields:

$$P(Z_i = k \mid \mathbf{Z}_{-(i)}, \mathbf{W}; \alpha, \beta) \propto (n_{m,-i}^{(k)} + \alpha_k) \frac{n_{k,-i}^{(t)} + \beta_t}{\sum_{t=1}^V n_{k,-i}^{(t)} + \beta_t} \quad (20)$$

where the counts $n_{\cdot,-i}^{(\cdot)}$ indicate that token i is excluded from the corresponding document or topic. Thus, sampling topic indexes using equation (20) for each word in a document and across documents until convergence allows us to approximate the posterior distribution given by (14). As noted in Heinrich (2009), the procedure itself uses only five larger data structures; the count variables $n_m^{(k)}$ and $n_k^{(t)}$, which have dimension $M \times K$ and $K \times V$, respectively, their row sums n_m and n_k , as well as the state variable $z_{m,n}$ with dimension W .

With one simulated sample of the posterior distribution for $P(\mathbf{Z} \mid \mathbf{W}; \alpha, \beta)$, φ and θ can be estimated from:

$$\hat{\varphi}_{k,t} = \frac{n_k^{(t)} + \beta_t}{\sum_{t=1}^V n_k^{(t)} + \beta_r} \quad (21)$$

and

$$\hat{\theta}_{m,k} = \frac{n_m^{(k)} + \alpha_k}{\sum_{k=1}^K n_m^{(k)} + \alpha_k} \quad (22)$$

In the analysis of the main paper the average of the estimated $\hat{\theta}$ and $\hat{\varphi}$ from the 10 last samples of the stored Gibbs simulations are used to construct the daily news topic frequencies.¹⁶ In un-reported experiments, the topic extraction results reported in Section 3.2 do not change much when choosing other samples for inference, for example using the last sample only.

The model is estimated using 7500×10 draws. The first 15000 draws of the sampler are disregarded, and only every 10th draw of the remaining simulations are recorded and used for inference. Because of the size of the regional data sets, see Section 3, we run into memory constraints if trying to use the whole cleaned corpus for estimation. For this reason we randomly sample, without replacement, up to 1.5 million articles from each data set.¹⁷ These samples are then used for estimating the word and topic distributions. However, when we construct daily topic frequencies, see Appendix C.2 below, all articles within each regional data set is used.

Before estimation three parameters need to be pre-defined: the number of topics and the two parameter vectors of the Dirichlet priors, α and β . Here, symmetric Dirichlet

¹⁶Because of lack of identifiability, the estimates of $\hat{\theta}$ and $\hat{\varphi}$ can not be combined across samples for an analysis that relies on the content of specific topics. However, statistics insensitive to permutation of the underlying topics can be computed by aggregating across samples, see Griffiths and Steyvers (2004).

¹⁷Note here that this step only applies to the US and euro area corpus, as the categorized data set for Japan is of a much smaller size already.

priors, with α and β each having a single value, are used. In turn, these are defined as a function of the number of topics and unique words:

$$\alpha = \frac{50}{K}, \text{ and } \beta = \frac{200}{N}$$

The choice of K is discussed in Section 3.2. In general, lower (higher) values for α and β will result in more (less) decisive topic associations. The values for the Dirichlet hyper-parameters also reflect a clear compromise between having few topics per document and having few words per topic. In essence, the prior specification used here is the same as the one advocated by Griffiths and Steyvers (2004).

C.2 Estimating daily topic frequencies

Using the posterior estimates from the LDA model, the frequency with which each topic is represented in the newspaper for a specific day is computed. This is done by first collapsing all the articles in the newspaper for one specific day into one document. Following Heinrich (2009) and Hansen et al. (2018), a procedure for querying documents outside the set on which the LDA is estimated is then implemented. In short, this corresponds to using the same Gibbs simulations as described above, but with the difference that the sampler is run with the estimated parameters $\Phi = \{\varphi_k\}_{k=1}^K$ and hyper-parameter α held constant.

Denote by $\tilde{W}$ the vector of words in the newly formed document. Topic assignments, $\tilde{Z}$, for this document can then be estimated by first initializing the algorithm by randomly assigning topics to words and then performing a number of Gibbs iterations using:

$$P(\tilde{Z}_i = k \mid \tilde{\mathbf{Z}}_{-(i)}, \tilde{\mathbf{W}}; \alpha, \beta) \propto (n_{\tilde{m}, -i}^{(k)} + \alpha_k) \hat{\varphi}_{k,t} \quad (23)$$

Since $\hat{\varphi}_{k,t}$ does not need to be estimated when sampling from (23), fewer iterations are needed to form the topic assignment index for the new document than when learning both the topic and word distributions. Here 2000 iterations are performed, and only the average of every 10th draw is used for the final inference. After sampling, the topic distribution can be estimated as before:

$$\tilde{\theta}_{\tilde{m}, k} = \frac{n_{\tilde{m}}^{(k)} + \alpha_k}{\sum_{k=1}^K n_{\tilde{m}}^{(k)} + \alpha_k} \quad (24)$$

C.3 News Topics as time series

Given knowledge of the topics (and their distributions), the topic decompositions are translated into time series. To do this, we proceed in three steps:

Step 1. For each day, the frequency with which each topic is represented in the newspaper that day is calculated. This is done by collapsing all the articles in the newspaper

for a particular day into one document, and then computing, using the estimated word distribution for each topic, the topic frequencies for this newly formed document. See Appendix C.2 for details. We label these time series $X_{t,k}$. By construction, across all topics, this number will sum to one for any given day. On average, across the whole sample, each topic will have a more or less equal probability of being represented in the newspaper. Across shorter time periods, i.e., days, the variation can be substantial.¹⁸

Step 2. Since the time series objects constructed in *Step 1* will be intensity measures, i.e., reflecting how much DN writes about a given topic at a specific point in time, their tone is not identified. That is, whether the news is positive or negative. To mediate this, a sign identified dataset based on the number of positive relative to negative words in the text is constructed. In particular, for each day t , all M_t newspaper articles that day, and each news topic, the article that news topic k describes the best is found. Given knowledge of this topic article mapping, positive/negative words in the articles are identified using an external word list and simple word counts. The word list used here is the *Harvard IV-4 Psychological Dictionary*.¹⁹

The count procedure delivers two statistics for each article, containing the number of positive and negative words. These statistics are then normalized such that each article observation reflects the fraction of positive and negative words, i.e.:

$$Post_{t,m_t} = \frac{\#positive\ words_{m_t}}{\#total\ words_{m_t}} \quad Neg_{t,m_t} = \frac{\#negative\ words_{m_t}}{\#total\ words_{m_t}} \quad (25)$$

The overall mood of article m_t , for $m_t = 1, \dots, M_t$ at day t , is defined as:

$$S_{t,m_t} = Post_{t,m_t} - Neg_{t,m_t} \quad (26)$$

Using the S_{t,m_t} statistic and the topic article mapping described above, the sign of each topic is adjusted as:

$$\tilde{X}_{t,k} = S_{t,m_t} X_{t,k}^{m_t}$$

where the m_t superscript is used on the topic frequency time series $X_{t,k}$ to highlight that topic k is mapped to article m_t .

Step 3. To remove daily noise from the topic time series in the $\tilde{X}_{t,k}$ dataset, each topic time series is filtered using a 60 day (backward looking) moving average filter. As is

¹⁸Note that the construction described in *Step 1* does not mean that only one topic is used as representative for a given day. For such an assumption, topic models other than the LDA would have been more appropriate.

¹⁹The word list can be obtained upon request. Counting the number of positive and negative words in a given text using the *Harvard IV-4 Psychological Dictionary* is a standard methodology in this branch of the literature, see, e.g., Tetlock et al. (2008). In finance, Loughran and McDonald (2011) among others, show that word lists developed for other disciplines mis-classify common words in financial text, and suggest an alternative (English language) list. We leave it for future research to investigate if this also holds for macroeconomic applications.

common in factor model studies, we also standardize the data prior to estimation (Stock and Watson (2012)).

Appendix D The Dynamic Factor model, estimation, and prediction

The mixed-frequency time-varying Dynamic Factor Model used for estimating the daily news-based coincident indexes builds on work in Thorsrud (2016b,a). A compact version of the model was described in Section 4. Below follows a more detailed description. First, the observation and transition equations of the system can be written as:

$$\mathbf{y}_t = \mathbf{Z}_t \mathbf{a}_t + \mathbf{e}_t \quad (27a)$$

$$\mathbf{a}_t = \mathbf{F}_t \mathbf{a}_{t-1} + \mathbf{R}_t \boldsymbol{\Sigma}_t \boldsymbol{\omega}_t \quad (27b)$$

$$\mathbf{e}_t = \mathbf{P} \mathbf{e}_{t-1} + \mathbf{u}_t \quad (27c)$$

with

$$\begin{aligned} \mathbf{y}_t &= \begin{pmatrix} \mathbf{y}_t^{k_q} \\ \mathbf{y}_t^{k_m} \\ \mathbf{y}_t^d \end{pmatrix} \quad \mathbf{Z}_t = \begin{pmatrix} \bar{\mathbf{Z}}^{k_q} & 0 & 0 \\ 0 & \bar{\mathbf{Z}}^{k_m} & 0 \\ 0 & 0 & \mathbf{z}_t^d \end{pmatrix} \quad \mathbf{a}_t = \begin{pmatrix} \mathbf{a}_t^{k_q} \\ \mathbf{a}_t^{k_m} \\ \mathbf{a}_t^d \end{pmatrix} \quad \mathbf{e}_t = \begin{pmatrix} \mathbf{e}_{1,t}^{k_q} \\ \mathbf{e}_{2,t}^{k_m} \\ \mathbf{e}_{3,t}^d \end{pmatrix} \\ \mathbf{F}_t &= \begin{pmatrix} \boldsymbol{\Upsilon}_t^{k_q} & 0 & -\boldsymbol{\pi}_t^{k_q} \Phi \\ 0 & \boldsymbol{\Upsilon}_t^{k_m} & -\boldsymbol{\pi}_t^{k_m} \Phi \\ 0 & 0 & \Phi \end{pmatrix} \quad \mathbf{R}_t = \begin{pmatrix} 1 & 0 & -\boldsymbol{\pi}_t^{k_q} \\ 0 & 1 & -\boldsymbol{\pi}_t^{k_m} \\ 0 & 0 & 1 \end{pmatrix} \quad \boldsymbol{\Sigma}_t = \begin{pmatrix} \boldsymbol{\sigma}_{t,\omega_q} & 0 & 0 \\ 0 & \boldsymbol{\sigma}_{t,\omega_m} & 0 \\ 0 & 0 & \boldsymbol{\sigma}_{t,\omega_d} \end{pmatrix} \\ \boldsymbol{\omega}_t &= \begin{pmatrix} \boldsymbol{\omega}_{t,q} \\ \boldsymbol{\omega}_{t,m} \\ \boldsymbol{\omega}_{t,d} \end{pmatrix} \quad \mathbf{P} = \begin{bmatrix} \Phi^{k_q} & 0 & 0 \\ 0 & \Phi^{k_m} & 0 \\ 0 & 0 & \Phi^d \end{bmatrix} \quad \mathbf{u}_t = \begin{bmatrix} \mathbf{u}_t^{k_q} \\ \mathbf{u}_t^{k_m} \\ \mathbf{u}_t^d \end{bmatrix} \end{aligned}$$

where t is the daily time index, k_q , k_m , and d denote the quarterly, monthly and daily observation intervals, respectively, and the model has been written with simple autoregressive time series processes of order one for notational simplicity.

The time-varying factor loadings are modeled as random walks following the Latent Threshold Model (LTM) idea introduced by Nakajima and West (2013). For example, for one particular element in the $\mathbf{z}_t^d$ vector, $z_{i,t}$, the LTM structure can be written as:

$$z_{i,t} = z_{i,t}^* \varsigma_{i,t} \quad \varsigma_{i,t} = I(|z_{i,t}^*| \geq d_i) \quad (28)$$

where

$$z_{i,t}^* = z_{i,t-1}^* + w_{i,t} \quad (29)$$

with $w_{i,t} \sim i.i.d.N(0, \sigma_{i,w}^2)$, and $\mathbf{w}_t \sim i.i.d.N(0, \mathbf{W})$ where $\mathbf{W}$ is a diagonal matrix. In (28) $\varsigma_{i,t}$ is a zero one variable, whose value depends on the indicator function $I(|z_{i,t}^*| \geq d_i)$. If $|z_{i,t}^*|$ is above the threshold value d_i , then $\varsigma_{i,t} = 1$, otherwise $\varsigma_{i,t} = 0$.

Stochastic volatility, stemming from $\mathbf{\Omega}_t = \mathbf{\Sigma}_t \mathbf{\Sigma}_t'$, is assumed to follow independent random walk processes:

$$\log(\boldsymbol{\sigma}_{t,\omega}) = \log(\boldsymbol{\sigma}_{t-1,\omega}) + \mathbf{b}_{t,\cdot}, \quad \mathbf{b}_{t,\cdot} \sim i.i.d.N(0, \mathbf{B}) \quad (30)$$

where $\mathbf{B}$ is a diagonal matrix.

Finally, the vectors of error terms, $\boldsymbol{\omega}_t$, $\mathbf{b}_t$, $\mathbf{u}_t$, and $\mathbf{w}_t$ are assumed to be mutually independent:

$$\begin{pmatrix} \boldsymbol{\omega}_t \\ \mathbf{b}_t \\ \mathbf{u}_t \\ \mathbf{w}_t \end{pmatrix} \sim i.i.d.N \left(\begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \begin{pmatrix} I & 0 & 0 & 0 \\ 0 & \mathbf{B} & 0 & 0 \\ 0 & 0 & \mathbf{U} & 0 \\ 0 & 0 & 0 & \mathbf{W} \end{pmatrix} \right)$$

The model's hyper-parameters are $\mathbf{B}$, $\mathbf{U}$, $\mathbf{W}$, $\mathbf{F}_t$, $\mathbf{P}$, and $\mathbf{d}$. Inside $\mathbf{F}_t$ and $\mathbf{R}_t$, the parameters $\boldsymbol{\pi}_t^k$ and $\boldsymbol{\Upsilon}_t^k$ are time-varying, but their evolution is deterministic and need not be estimated, confer Appendix D.7. Thus, the only time-varying parameters to be estimated are those in $\mathbf{Z}_t$ and $\mathbf{\Sigma}_t$, which together with $\mathbf{a}_t$, are the model's unobserved state variables.

Estimation consists of sequentially drawing the model's unobserved state variables and hyper-parameters utilizing 5 blocks until convergence is achieved. In essence, each block involves exploiting the state space nature of the model using the Kalman Filter and the simulation smoother suggested by Carter and Kohn (1994), coupled with a Metropolis-Hastings step to simulate the time-varying loadings. Below we describe each block in greater detail. Our main results are obtained from 50000 iterations. The first 10000 are discarded and only every 10th of the remaining are used for inference.

For future reference and notational simplicity it will prove useful to define the following: $\mathbf{Y} = [\mathbf{y}_1, \dots, \mathbf{y}_T]'$, $\mathbf{A} = [\mathbf{a}_1, \dots, \mathbf{a}_T]'$, $\mathbf{Z} = [\mathbf{Z}_1, \dots, \mathbf{Z}_T]'$, $\mathbf{E} = [\mathbf{e}_1, \dots, \mathbf{e}_T]'$, $\mathbf{F} = [\mathbf{F}_1, \dots, \mathbf{F}_T]'$, and $\mathbf{Q} = [\mathbf{\Omega}_1, \dots, \mathbf{\Omega}_T]$.

D.1 Block 1: $\mathbf{A}|\mathbf{Y}, \mathbf{Z}, \mathbf{F}, \mathbf{P}, \mathbf{U}, \mathbf{Q}$

Equations (27a) and (27b) constitute a state space system we can use to draw the unobserved state $\mathbf{a}_t$ using the Carter and Kohn's multimove Gibbs sampling approach. However, to do so we need to make the errors in the observation equation conditionally i.i.d. Given knowledge of equation (27c), we can define $\mathbf{P}(L) = (\mathbf{I} - \mathbf{P}L)$ and pre-multiply

equation (27a) by $\mathbf{P}(L)$ to obtain the system:

$$\tilde{\mathbf{y}}_t = \tilde{\mathbf{Z}}_t \mathbf{a}_t + \mathbf{u}_t \quad (31a)$$

$$\mathbf{a}_t = \mathbf{F}_t \mathbf{a}_{t-1} + \mathbf{R}_t \boldsymbol{\Sigma}_t \boldsymbol{\omega}_t \quad (31b)$$

where $\tilde{\mathbf{y}}_t = \mathbf{P}(L)\mathbf{y}_t$ and $\tilde{\mathbf{Z}}_t = \mathbf{P}(L)\mathbf{Z}_t$.

Since all hyper-parameters and state variables, less $\mathbf{A}$, are known (or conditionally known), we can use the equations in (31) together with Carter and Kohn's multimove Gibbs sampling approach, see Appendix (E), to sample $\mathbf{a}_t$ from:

$$\mathbf{a}_T | \dots \sim N(\mathbf{a}_{T|T}, \mathbf{P}_{T|T}^a) \quad t = T \quad (32a)$$

$$\mathbf{a}_t | \dots \sim N(\mathbf{a}_{t|t, a_{t+1}}, \mathbf{P}_{t|t, a_{t+1}}^a) \quad t = T-1, T-2, \dots, 0 \quad (32b)$$

to get $\mathbf{A}$. Note here that the Kalman Filter can be run straightforwardly despite the fact that the $\tilde{\mathbf{y}}_t$ vector contains missing values, see Harvey (1990) for details.

D.2 Block 2: $\mathbf{Z}, d|Y, \mathbf{A}, \mathbf{P}, U, \mathbf{W}$ and $\mathbf{W}|\mathbf{Z}$

Conditionally on $\mathbf{A}$ the errors in (27a) are independent across the N variables in $\mathbf{y}_t$. Moreover, we have assumed that the covariance matrix $\mathbf{W}$ associated with the time-varying factor loadings in equation (29) is diagonal. Consequently, one can draw $\mathbf{Z}$ one equation at a time. As above, we deal with the fact that the errors in the observation equation are not conditionally i.i.d. by applying the quasi differencing operator, $\mathbf{P}(L)$, to each equation. Thus, for each i in N_d , we obtain the following Gaussian system:

$$\tilde{y}_{i,t}^j = \tilde{a}_t^j z_{i,t}^j + u_{i,t}^j \quad (33a)$$

$$z_{i,t}^j = z_{i,t}^* \varsigma_{i,t} \quad \varsigma_{i,t} = I(|z_{i,t}^*| \geq d_i) \quad (33b)$$

$$z_{i,t}^* = z_{i,t-1}^* + w_{i,t} \quad (33c)$$

where $\tilde{y}_{i,t}^j = (I - \Phi_i^j L)y_{i,t}^j$ and $\tilde{a}_t^j = (I - \Phi_i^j L)a_t^j$, for $j = k_q, k_m$, or d depending on the observation frequency of variable i .

To simulate from the conditional posterior of $z_{i,t}^*$ and d_i in (33), the procedure outlined in Nakajima and West (2013) is followed. That is, conditional on all the data and hyper-parameters, we draw the conditional posterior of $z_{i,t}^*$ sequentially for $t = 1 : T$, or $t = k, 2k, \dots$, for variables observed at a lower frequency than daily, using a Metropolis-Hastings (MH) sampler. As described in Nakajima and West (2013), the MH proposals come from a non-thresholded version of the model specific to each time t , or observation interval, as follows: Fixing $\varsigma_{i,t} = 1$, and dropping the j superscript for notational simplicity, take proposal distribution $N(z_{i,t}^* | m_t, M_t)$ where:

$$M_t^{-1} = \sigma_{i,u}^{-2} \tilde{a}_t \tilde{a}_t + \sigma_{i,w}^{-2} (I + 1) \quad (34a)$$

$$m_t = M_t [\sigma_{i,u}^{-2} \tilde{a}_t \tilde{y}_{i,t} + \sigma_{i,w}^{-2} \{(z_{i,t-1}^* + z_{i,t+1}^*) + (I - 1)z_{i,0}^*\}] \quad (34b)$$

for $t = 2 : T - 1$. For $t = 1$ and $t = T$, a slight modification is needed. Details can be found in Nakajima and West (2013). The candidate is accepted with probability:

$$\alpha(z_{i,t}^*, z_{i,t}^{p*}) = \min \left\{ 1, \frac{N(\tilde{y}_{i,t} | \tilde{a}_t z_{i,t}^p, \sigma_{i,u}^2) N(z_{i,t}^* | m_t, M_t)}{N(\tilde{y}_{i,t} | \tilde{a}_t z_{i,t}, \sigma_{i,u}^2) N(z_{i,t}^{p*} | m_t, M_t)} \right\} \quad (35)$$

where $z_{i,t} = z_{i,t}^* \varsigma_{i,t}$ is the current state, and $z_{i,t}^p = z_{i,t}^{p*} \varsigma_{i,t}^p$ is the candidate.

The independent latent thresholds in d_i can then be sampled conditional on the data and the hyper-parameters. For this, a direct MH algorithm is employed. Let $d_{i,-j} = d_{i,0:s} \setminus d_{i,j}$. A candidate is drawn from the current conditional prior, $d_{i,j}^p \sim U(0, |\beta_0| + K)$, where K is described below, and accepted with probability:

$$\alpha(d_{i,j}, d_{i,j}^p) = \min \left\{ 1, \prod_{t=1}^T \frac{N(\tilde{y}_{i,t} | \tilde{a}_t z_{i,t}^p, \sigma_{i,u}^2)}{N(\tilde{y}_{i,t} | \tilde{a}_t z_{i,t}, \sigma_{i,u}^2)} \right\} \quad (36)$$

where $z_{i,t}$ is the state based on the current thresholds $(d_{i,j}, d_{i,-j})$, and $z_{i,t}^p$ the candidate based on $(d_{i,j}^p, d_{i,-j})$.

Lastly, conditional on the data, the hyper-parameters and the time-varying parameters, we can sample the elements of $\mathbf{W}$ using the inverse Gamma distribution. Letting letters denoted with an underscore reflect the prior, then:

$$\sigma_{i,w}^2 | \dots \sim IG(\bar{v}^w, \bar{\sigma}_{i,w}^2) \quad (37)$$

where $\bar{v}^w = T + \mathcal{T}^w$ and $\bar{\sigma}_{i,w}^2 = [\sigma_{i,w}^2 \mathcal{T}^w + \sum_{t=1}^T (z_{i,t}^* - z_{i,t-1}^*)'(z_{i,t}^* - z_{i,t-1}^*)] / \bar{v}^w$.

Notice here that the identifying restrictions, confer Section 4, put a restriction on the first element in the $N_d \times 1$ vector of daily observables. For this particular i , $z_{i,t} = z_{i,t}^* = 1$ for all t , and $\sigma_{i,w}^2 = 0$ and $d_i = 0$. Moreover, in the cases where $z_{i,t}^j = z_i^j$ for all time periods, i.e., static, inference becomes much simpler. This applies to $z_i^{k_q}$ and $z_i^{k_m}$ in all model specifications, but only to z_i^d in the model labeled NCI^{notvp} . Thus, after doing the transformation in (33a), the Normal-Gamma prior implies that:

$$z_i^j | \dots \sim N(\bar{z}_i^j, \bar{V}^{z_i^j}) \quad (38)$$

with

$$\bar{V}^{z_i^j} = (\mathcal{V}^{z_i^j} + \sum_{t=1}^T \tilde{a}_t^{j'} \sigma_{i,w}^{-2} \tilde{a}_t^j) \quad (39)$$

$$\bar{z}_i^j = \bar{V}^{z_i^j} (\mathcal{V}^{z_i^j} z_i^j + \sum_{t=1}^T \tilde{a}_t^{j'} \sigma_{i,w}^{-2} \tilde{y}_{i,t}^j) \quad (40)$$

and conditional on $\bar{z}_i^j, \sigma_{i,w}^2$ can be sampled from (37) with $\bar{\sigma}_{i,w}^2 = [\sigma_{i,w}^2 \mathcal{T}^w + \sum_{t=1}^T (\tilde{y}_{i,t}^j - \tilde{a}_t^j z_{i,t}^j)'(\tilde{y}_{i,t}^j - \tilde{a}_t^j z_{i,t}^j)] / \bar{v}^w$.

D.3 Block 3: $U|Y, A, P$ and $P|Y, A, U$

Conditional on Y , A , and P we can use $\tilde{y}_{i,t}^j = (I - \Phi_i^j L)y_{i,t}^j$ and $\tilde{a}_t^j = (I - \Phi_i^j L)a_t^j$ defined above, and simulate the errors in U from the inverse Gamma distribution:

$$\sigma_{i,u}^2 | \dots \sim IG(\bar{v}^u, \bar{\sigma}_{i,u}^2) \quad (41)$$

where $\bar{v}^u = T + \underline{T}^u$, $\bar{\sigma}_{i,u}^2 = [\sigma_{i,u}^2 \underline{T}^u + \sum_{t=1}^T (\tilde{y}_{i,t} - \tilde{a}_t z_{i,t})'(\tilde{y}_{i,t} - \tilde{a}_t z_{i,t})] / \bar{v}^u$, and the superscripts j are dropped for notational simplicity.

Given U , Y , and A , it follows that each element of E is given by:

$$e_{i,t} = y_{i,t} - z_{i,t} a_t \quad (42)$$

From this we can then sample the Φ elements of P using the standard independent Normal-Gamma prior. Accordingly, for each non-restricted element in P :

$$\Phi_i | \dots \sim N(\bar{\Phi}_i, \bar{V}_i^\Phi)_{I[s(\Phi_i)]} \quad (43)$$

with

$$\bar{V}_i^\Phi = (V_i^{\Phi^{-1}} + \sum_{t=1}^T e'_{i,t-1} \sigma_{i,u}^{-2} e_{i,t-1})^{-1} \quad (44)$$

$$\bar{\Phi}_i = \bar{V}_i^\Phi (V_i^{\Phi^{-1}} \Phi_i + \sum_{t=1}^T e'_{i,t-1} \sigma_{i,u}^{-2} e_{i,t}) \quad (45)$$

and $I[s(\Phi_i)]$ is an indicator function used to denote that the roots of Φ_i lie outside the unit circle.

D.4 Block 4: $F|A, \Omega$

Conditional on A , the transition equation in (27b) is independent of the rest of the model. Moreover, conditional on knowing Ω , and with the restriction that $\Sigma_t = \sigma_{t,\omega_d}$, all elements in F_t and R_t are known except Φ . Thus, we can focus on the last element in a_t (a_t^d), and draw Φ using the independent Normal-Gamma prior. Continuing with letting letters denoted with an underscore reflect the prior, the conditional posterior of Φ is:

$$\Phi | \dots \sim N(\bar{\Phi}, \bar{V}^\Phi)_{I[s(\Phi)]} \quad (46)$$

with

$$\bar{V}^\Phi = (V^{\Phi^{-1}} + \sum_{t=1}^T a_{t-1}^{d'} \sigma_{t,\omega_d}^{-2} a_{t-1}^d)^{-1} \quad (47)$$

$$\bar{\Phi} = \bar{V}^\Phi (V^{\Phi^{-1}} \Phi + \sum_{t=1}^T a_{t-1}^{d'} \sigma_{t,\omega_d}^{-2} a_t^d) \quad (48)$$

and $I[s(\Phi)]$ is an indicator function used to denote that the roots of Φ lie outside the unit circle.

D.5 Block 5: $\Omega|F, A, B$, and $B|\Omega$

Conditional on the elements a_t^d and Φ of A and F , we can define $\hat{a}_t^d = a_t^d - \Phi a_{t-1}^d$, and write the last line of equation (27b) as:

$$\hat{a}_t^d = \sigma_{t,\omega_d} \omega_{t,d} \quad (49)$$

Together with the transition equation in (30), the observation equation in (49) constitutes a nonlinear state space system. The nonlinearity can be converted into a linear one by squaring and taking logarithms of every element of (49), yielding:

$$\hat{a}_t^{d*} = 2h_t^\sigma + \omega_{t,d}^* \quad (50a)$$

$$h_t^\sigma = h_{t-1}^\sigma + b_{t,d} \quad (50b)$$

where $h_t^\sigma = \log(\sigma_{t,\omega_d})$, $\omega_{t,d}^* = \log(\omega_{t,d}^2)$, $\hat{a}_t^{d*} = \log((\hat{a}_t^d)^2 + \bar{c})$, and $\bar{c} = 0.001$ is an offsetting constant added to the latter expression to avoid potentially taking the log of zero.

Now, the system in (50) is linear, but it has a non-Gaussian state space form, because the innovations in the observation equation are distributed as $\log \chi^2(1)$. In order to further transform the system into a Gaussian one, a mixture of normals approximation of the $\log \chi^2(1)$ distribution is used. Following Kim et al. (1998), we select a mixture of seven normal densities with component probabilities q_γ , mean $m_\gamma - 1.2704$, and variances v_γ^2 , for $\gamma = 1, \dots, 7$. The constants $q_\gamma, m_\gamma, v_\gamma^2$ are chosen to match a number of moments of the $\log \chi^2(1)$ distribution. Accordingly, conditionally on $\hat{a}_t^{d*}$ and h_t^σ , we can sample a selection matrix $\tilde{s}_T = [s_1, \dots, s_T]'$ as:

$$Pr(s_{l,t} = \gamma | \hat{a}_t^{d*}, h_t^\sigma) \propto q_\gamma f_N(\hat{a}_t^{d*} | 2h_t^\sigma + m_\gamma - 1.2704, v_\gamma^2) \quad \gamma = 1, \dots, 7 \quad l = 1, \dots, q \quad (51)$$

and use the selection matrix to select which member of the mixture of the normal approximations that should be used to construct the covariance matrix of $\omega_{t,d}^*$ and adjust the mean of $\hat{a}_t^{d*}$ at every point in time. In turn, conditional on B , these adjusted terms are used to recursively recover h_t^σ , for $t = 1, \dots, T$ using the Carter and Kohn's multimove Gibbs sampling approach (Appendix (E)):

$$h_T^\sigma | \dots \sim N(h_{T|T}^\sigma, P_{T|T}^{h^\sigma}), \quad t = T \quad (52a)$$

$$h_t^\sigma | \dots \sim N(h_{t|t, h_{t+1}^\sigma}^\sigma, P_{t|t, h_{t+1}^\sigma}^{h^\sigma}), \quad t = T-1, T-2, \dots, 0 \quad (52b)$$

Finally, conditional on h_t^σ , the posterior of $B = \sigma_{b_d}^2$ is drawn from the inverse Gamma distribution:

$$\sigma_{b_d}^2 | \dots \sim IG(\bar{v}^{b_d}, \bar{\sigma}_{b_d}^2) \quad (53)$$

where $\bar{v}^{b_d} = T + \underline{T}^{b_d}$, $\bar{\sigma}_{b_d}^2 = [\underline{\sigma}_{b_d}^2 \underline{T}^{b_d} + \sum_{t=1}^T (h_t^\sigma - h_{t-1}^\sigma)'(h_t^\sigma - h_{t-1}^\sigma)] / \bar{v}^{b_d}$.

D.6 Priors

To implement the MCMC algorithm, prior specifications for the initial state variables $\mathbf{a}_0$, $\mathbf{Z}_0$, $\mathbf{\Sigma}_0$, and for the hyper-parameters $\mathbf{B}$, $\mathbf{U}$, $\mathbf{W}$, $\mathbf{F}_t$, $\mathbf{P}$, and $\mathbf{d}$ are needed. The prior specifications used for the initial states take the following form: $\mathbf{a}_0 \sim N(0, I \cdot 10)$, $\mathbf{Z}_0 \sim N(0, I)$, and $\mathbf{\Sigma}_0 \sim N(1, I)$. The priors for the hyper-parameters Φ and Φ_i , which are part of the $\mathbf{F}_t$ and $\mathbf{P}$ matrices, respectively, are set to $\Phi \sim N(0, I)$ and $\Phi_i \sim N(0, 0.5)$. For the constant parameters in $\mathbf{Z}_t$, i.e., $\mathbf{Z}^k$, we assume for each element i that $z_i^k \sim N(1, 1)$. The priors for $\mathbf{B}$, $\mathbf{U}$, and $\mathbf{W}$, are all from the Inverse-Gamma distribution, where the first element in each prior distribution is the shape parameter, and the second the scale parameter: $\sigma_{b_d}^2 \sim IG(\underline{T}^{b_d}, \kappa_{b_d}^2)$ with $\underline{T}^{b_d} = T \cdot 0.1$ and $\kappa_{b_d} = 0.01$; $\sigma_{i,u}^2 \sim IG(\underline{T}^u, \kappa_u^2)$ with $\underline{T}^u = T \cdot 0.5$ and $\kappa_u = 0.3$; $\sigma_{i,w}^2 \sim IG(\underline{T}^w, \kappa_w^2)$ with $\underline{T}^w = T \cdot 1$ and $\kappa_w = 0.003$, where T is the sample size. In sum, as the full sample contains up to 9000 observations, these priors are informative for the variance terms associated with the time-varying factor loadings, but less so for the other parameters. To draw the latent threshold, $\mathbf{d}$, a tuning parameter controlling our prior belief concerning the marginal sparsity probability needs to be defined. A neutral prior will support a range of sparsity values in order to allow the data to inform on relevant values. Here we set it to 0.4, which according to the analysis in [Thorsrud \(2016b\)](#) provides a reasonable prior in terms of balancing the degree of sparsity and potential over-fitting.

Finally, the MCMC simulations are initialized using simple OLS estimates obtained using the cross-sectional mean of the news topics as a measure of the daily business cycle index.

D.7 The cumulator variable approach

As is common in mixed-frequency models, lower frequency variables are treated as daily series with missing observations ([Forni and Marcellino \(2013\)](#)), and time aggregation from higher to lower frequency is restricted as follows for a generic variable y_t^k :

$$\begin{aligned} y_t^k &= \log(v_{1,t}^k) - \log(v_{1,t-k}^k) \approx \log\left(\sum_{i=0}^{k-1} v_{1,t-i}\right) - \log\left(\sum_{i=k}^{2k-1} v_{1,t-i}\right) \\ &\approx \sum_{i=0}^{k-1} \log(v_{1,t-i}) - \sum_{i=k}^{2k-1} \log(v_{1,t-i}) = \sum_{i=0}^{2k-2} \omega_i^k y_{1,t-i}, \quad t = k, 2k, \dots \end{aligned} \quad (54)$$

where y_t^k is the observed low frequency growth rate, v_t^k its level, and $\omega_i^k = i + 1$ for $i = 0, \dots, k - 1$; $\omega_i^k = 2k - i - 1$ for $i = k, \dots, 2k - 2$; and $\omega_i^k = 0$ otherwise. Imposing a common factor structure for y_t^k , it follows from (54) that at the observation interval:

$$y_t^k = \sum_{i=0}^{2k-2} \omega_i^k y_{1,t-i} = \sum_{i=0}^{2k-2} \omega_i^k (z a_{t-i}^d + e_{t-i}) \quad (55)$$

A caveat with the model formulation in (55) is that it increases the number of state variables in the system considerably. For example, when aggregation is from daily to quarterly frequency, the number of elements in the state vector exceed 180, posing significant challenges for estimation.²⁰ To limit the size of the state vector, temporal aggregation is handled using a double cumulator variable approach as in Banbura et al. (2013). The temporal aggregator variables are recursively updated such that at the end of each respective period we have:

$$a_t^k = \sum_{i=0}^{2k-2} \omega_i^k a_{t-i}, \quad t = k, 2k, \dots, \quad (56)$$

As shown below, these recursions can be computed with the help of only two additional state variables and selection and weight matrices. In (27a) this is reflected in the partition $\mathbf{a}_t^k = \begin{pmatrix} a_t^k & \bar{a}_t^k \end{pmatrix}'$, the selection matrix $\mathbf{\Upsilon}_t^k$, and the vector $\boldsymbol{\pi}_t^k$ which contains the aggregation weights ω_i^k . Accordingly, $\bar{\mathbf{Z}}^k = \begin{pmatrix} \mathbf{Z}^k & \mathbf{0} \end{pmatrix}$. Notice here that the factor loadings are static. Allowing for time-varying loadings for the low frequency variables will be in conflict with the aggregation scheme in (55) and (56).

The time aggregation structure of the model, given by equation (55), introduces moving average terms into the idiosyncratic errors for the monthly and quarterly variables. In the case of only one monthly and quarterly variable this is captured by the $\mathbf{R}_t \boldsymbol{\Sigma}_t \boldsymbol{\omega}_t$ term in (1b). However, allowing for such time series patterns, we find that the model becomes substantially more difficult to estimate. For this reason we follow the specification adopted in Banbura et al. (2013), and assume *i.i.d.* errors at the monthly and quarterly observation intervals. This amounts to restricting $\mathbf{R}_t = \begin{bmatrix} -\boldsymbol{\pi}_t^{k_q} & -\boldsymbol{\pi}_t^{k_m} & 1 \end{bmatrix}'$, $\boldsymbol{\Sigma}_t = \sigma_{t,\omega_d}$, $\boldsymbol{\omega}_t = \omega_{t,d}$, and $\boldsymbol{\Phi}^{k_q} = \boldsymbol{\Phi}^{k_m} = \mathbf{0}$.

From equation (56) we had that:

$$a_t^k = \sum_{i=0}^{2k-2} \omega_i^k a_{t-i}, \quad t = k, 2k, \dots, \quad (57)$$

As shown in Banbura et al. (2013), this expression can be computed recursively with the help of two (additional) state variables. In particular, by introducing the auxiliary variable $\bar{a}_t^k$, a_t^k is obtained recursively as follows:

$$\mathbf{a}_t^k = \begin{pmatrix} a_t^k \\ \bar{a}_t^k \end{pmatrix} = \begin{cases} \begin{pmatrix} \bar{a}_{t-1}^k + \omega_{k-1}^k a_t \\ 0 \end{pmatrix}, & t = 1, k+1, 2k+1, \dots, \\ \begin{pmatrix} a_{t-1}^k + \omega_{R(k-t,k)}^k a_t \\ \bar{a}_{t-1}^k + \omega_{R(k-t,k)+k}^k a_t \end{pmatrix}, & \text{otherwise,} \end{cases} \quad (58)$$

²⁰In a constant parameter setting, Aruoba et al. (2009) employ Maximum Likelihood estimation where one evaluation of the likelihood takes roughly 20 seconds. As Bayesian estimation using MCMC requires a large number of iterations, the problem is infeasible in terms of computation time.

where $R(\cdot, k)$ denotes the positive remainder of the division by k . In turn, the expressions in (58) can be implemented in the time-varying mixed frequency DFM with the following weight vector $\boldsymbol{\pi}_t^k$ and selection matrix $\boldsymbol{\Upsilon}_t^k$:

$$\boldsymbol{\pi}_t^k = \begin{cases} \begin{pmatrix} -\omega_{k-1}^k \\ 0 \end{pmatrix}, & t = 1, k+1, \dots, \\ \begin{pmatrix} -\omega_{R(k-t,k)}^k \\ -\omega_{R(k-t,k)+k}^k \end{pmatrix}, & \text{otherwise,} \end{cases} \quad \boldsymbol{\Upsilon}_t^k = \begin{cases} \begin{pmatrix} 0 & 1 \\ 0 & 0 \end{pmatrix}, & t = 1, k+1, \dots, \\ I_2, & \text{otherwise,} \end{cases} \quad (59)$$

Generally, the mixed frequency framework described by equations (57), (58), and (59) can handle temporal aggregation from higher to lower frequencies for a range of k values. In the model formulation described in Section 4, only $k = k_q$ is considered, where the k 's refer to the (average) number of days in a quarter.

To deal with different number of days per quarter, a small adjustment needs to be implemented. Here we follow Banbura et al. (2013) and make the approximation that:

$$v_t^k = \frac{k}{k_t} \sum_{i=0}^{k_t-1} v_{t-i}, \quad t = k_1, k_1 + k_{k_1+1}, \dots \quad (60)$$

where k_t is the number of business days in the period (month or quarter) that contains day t and k is the average number of business days per period over the sample. As shown in Banbura et al. (2013), this results in time-varying weights, and the formulas above should be updated with: $\omega_{t,i}^k = k \frac{i+1}{k_t}$ for $i = 0, 1, \dots, k_t - 1$; $\omega_{t,i}^k = k \frac{k_t + k_t - k_t - i - 1}{k_t - k_t}$ for $i = k_t, k_t + 1, \dots, k_t + k_t - k_t - 2$; and $\omega_{t,i}^k = 0$ otherwise.

Appendix E The Carter and Kohn algorithm

Consider a generic state space system, written in companion form, and described by:

$$\mathbf{y}_t = \mathbf{Z}\mathbf{a}_t + \mathbf{u}_t \sim N(0, \mathbf{U}) \quad (61a)$$

$$\mathbf{a}_t = \mathbf{F}\mathbf{a}_{t-1} + \boldsymbol{\omega}_t \sim N(0, \mathbf{Q}) \quad (61b)$$

where the parameters are assumed to be known and constant for notational simplicity, and we wish to estimate the latent state $\mathbf{a}_t$ for all $t = 1, \dots, T$. To do so, we apply Carter and Kohn's multimove Gibbs sampling approach (Carter and Kohn (1994)).

First, because the state space model given in equation (61) is linear and (conditionally) Gaussian, the distribution of $\mathbf{a}_t$ given $\mathbf{Y}$ and that of $\mathbf{a}_t$ given $\mathbf{a}_{t+1}$ and $\mathbf{Y}$ for $t = T - 1, \dots, 1$ are also Gaussian:

$$\mathbf{a}_T | \mathbf{Y} \sim N(\mathbf{a}_{T|T}, \mathbf{P}_{T|T}), \quad t = T \quad (62a)$$

$$\mathbf{a}_t | \mathbf{Y}, \mathbf{a}_{t+1} \sim N(\mathbf{a}_{t|t, \mathbf{a}_{t+1}}, \mathbf{P}_{t|t, \mathbf{a}_{t+1}}), \quad t = T - 1, T - 2, \dots, 1 \quad (62b)$$

where

$$\mathbf{a}_{T|T} = E(\mathbf{a}_T|\mathbf{Y}) \quad (63a)$$

$$\mathbf{P}_{T|T} = Cov(\mathbf{a}_T|\mathbf{Y}) \quad (63b)$$

$$\mathbf{a}_{t|t, \mathbf{a}_{t+1}} = E(\mathbf{a}_t|\mathbf{Y}, \mathbf{a}_{t+1}) = E(\mathbf{a}_t|\mathbf{a}_{t|t}, \mathbf{a}_{t|t+1}) \quad (63c)$$

$$\mathbf{P}_{t|t, \mathbf{a}_{t+1}} = Cov(\mathbf{a}_t|\mathbf{Y}, \mathbf{a}_{t+1}) = Cov(\mathbf{a}_t|\mathbf{a}_{t|t}, \mathbf{a}_{t|t+1}) \quad (63d)$$

Given $\mathbf{a}_{0|0}$ and $\mathbf{P}_{0|0}$, the unknown states $\mathbf{a}_{T|T}$ and $\mathbf{P}_{T|T}$ needed to draw from (62a) can be estimated from the (conditionally) Gaussian Kalman Filter as:

$$\mathbf{a}_{t|t-1} = \mathbf{F}\mathbf{a}_{t-1|t-1} \quad (64a)$$

$$\mathbf{P}_{t|t-1} = \mathbf{F}\mathbf{P}_{t-1|t-1}\mathbf{F}' + \mathbf{Q} \quad (64b)$$

$$\mathbf{K}_t = \mathbf{P}_{t|t-1}\mathbf{Z}'(\mathbf{Z}\mathbf{P}_{t|t-1}\mathbf{Z}' + \mathbf{U})^{-1} \quad (64c)$$

$$\mathbf{a}_{t|t} = \mathbf{a}_{t|t-1} + \mathbf{K}_t(\mathbf{y}_t - \mathbf{Z}\mathbf{a}_{t|t-1}) \quad (64d)$$

$$\mathbf{P}_{t|t} = \mathbf{P}_{t|t-1} - \mathbf{K}_t\mathbf{Z}\mathbf{P}_{t|t-1} \quad (64e)$$

At $t = T$, equation 64d and 64e above, together with equation 62a, can be used to draw $\mathbf{a}_{T|T}$. $\mathbf{a}_{t|t, \mathbf{a}_{t+1}}$, for $t = T - 1, T - 2, \dots, 1$, can then be simulated based on 62b, where $\mathbf{a}_{t|t, \mathbf{a}_{t+1}}$ and $\mathbf{P}_{t|t, \mathbf{a}_{t+1}}$ are generated from the following updating equations:

$$\mathbf{a}_{t|t, \mathbf{a}_{t+1}} = \mathbf{a}_{t|t} + \mathbf{P}_{t|t}\mathbf{F}'(\mathbf{F}\mathbf{P}_{t|t}\mathbf{F}' + \mathbf{Q})^{-1}(\mathbf{a}_{t+1} - \mathbf{F}\mathbf{a}_{t|t}) \quad (65a)$$

$$\mathbf{P}_{t|t, \mathbf{a}_{t+1}} = \mathbf{P}_{t|t} + \mathbf{P}_{t|t}\mathbf{F}'(\mathbf{F}\mathbf{P}_{t|t}\mathbf{F}' + \mathbf{Q})^{-1}\mathbf{F}\mathbf{P}_{t|t} \quad (65b)$$

Appendix F Convergence of the Markov Chain Monte Carlo Algorithm

Table 17 summarizes the main convergence statistics used to check that the Gibbs sampler mixes well. In the first row of the table the mean, as well as the minimum and maximum, of the 10th-order sample autocorrelation of the posterior draws is reported. A low value indicates that the draws are close to independent. The second row of the table reports the relative numerical efficiency measure (RNE), proposed by Geweke (1992). The RNE measure provides an indication of the number of draws that would be required to produce the same numerical accuracy if the draws represented had been made from an i.i.d. sample drawn directly from the posterior distribution. An RNE value close to or below unity is regarded as satisfactory. Autocorrelation in the draws is controlled for by employing a 4 percent tapering of the spectral window used in the computation of the RNE.

As can be seen from the results reported in Table 17, the sampler seems to have converged. That is, the mean autocorrelations are all very close to zero, and the minimum

Table 17. Convergence statistics. The *AutoCorr* row reports the 10th-order sample autocorrelation of the draws, while the *RNE* row reports the relative numerical efficiency measure, proposed by Geweke (1992). For each entry we report the mean value together with the minimum and maximum value obtained across all parameters in parentheses.

Statistic	Parameters					
	U	B	P	F_t	W	d
Panel A: <i>NCI-US</i>						
AutoCorr	$\frac{-0.0}{(-0.1, 0.1)}$	$\frac{0.4}{(0.4, 0.4)}$	$\frac{-0.0}{(-0.1, 0.1)}$	$\frac{0.0}{(-0.1, 0.1)}$	$\frac{0.3}{(-0.0, 0.6)}$	$\frac{0.0}{(-0.1, 0.2)}$
RNE	$\frac{1.1}{(0.6, 2.0)}$	$\frac{0.1}{(0.1, 0.1)}$	$\frac{1.1}{(0.6, 1.7)}$	$\frac{1.2}{(0.8, 1.5)}$	$\frac{0.1}{(0.1, 0.5)}$	$\frac{0.8}{(0.1, 1.7)}$
Panel B: <i>NCI-Japan</i>						
AutoCorr	$\frac{-0.0}{(-0.1, 0.1)}$	$\frac{0.4}{(0.4, 0.4)}$	$\frac{0.0}{(-0.1, 0.1)}$	$\frac{-0.0}{(-0.0, 0.0)}$	$\frac{0.2}{(0.0, 0.5)}$	$\frac{0.0}{(-0.1, 0.2)}$
RNE	$\frac{1.1}{(0.7, 2.0)}$	$\frac{0.1}{(0.1, 0.1)}$	$\frac{1.1}{(0.7, 2.2)}$	$\frac{1.3}{(0.7, 1.6)}$	$\frac{0.2}{(0.1, 0.5)}$	$\frac{0.6}{(0.1, 1.5)}$
Panel C: <i>NCI-Euro</i>						
AutoCorr	$\frac{0.0}{(-0.1, 0.1)}$	$\frac{0.4}{(0.4, 0.4)}$	$\frac{0.0}{(-0.1, 0.1)}$	$\frac{-0.0}{(-0.1, 0.1)}$	$\frac{0.3}{(0.1, 0.6)}$	$\frac{0.0}{(-0.1, 0.2)}$
RNE	$\frac{1.1}{(0.6, 1.9)}$	$\frac{0.1}{(0.1, 0.1)}$	$\frac{1.1}{(0.5, 1.9)}$	$\frac{1.2}{(0.8, 1.7)}$	$\frac{0.2}{(0.1, 0.5)}$	$\frac{0.8}{(0.2, 1.6)}$

or maximum values obtained seldom exceed 0.1 in absolute value. Moreover, the mean RNE statistic does not exceed unity by a large margin for any of the parameters.