

Using noise filtering and sufficient dimension reduction method on unstructured economic data

Jae Keun Yoo^a, Yujin Park^a, Beomseok Seo^{1,b}

^aDepartment of Statistics, Ewha Womans University;

^bOffice of Economic Modeling and Policy Analysis, Bank of Korea

Abstract

Text indicators are increasingly valuable in economic forecasting, but are often hindered by noise and high dimensionality. This study aims to explore post-processing techniques, specifically noise filtering and dimensionality reduction, to normalize text indicators and enhance their utility through empirical analysis. Predictive target variables for the empirical analysis include monthly leading index cyclical variations, BSI (business survey index) All industry sales performance, BSI All industry sales outlook, as well as quarterly real GDP SA (seasonally adjusted) growth rate and real GDP YoY (year-on-year) growth rate. This study explores the Hodrick and Prescott filter, which is widely used in econometrics for noise filtering, and employs sufficient dimension reduction, a nonparametric dimensionality reduction methodology, in conjunction with unstructured text data. The analysis results reveal that noise filtering of text indicators significantly improves predictive accuracy for both monthly and quarterly variables, particularly when the dataset is large. Moreover, this study demonstrated that applying dimensionality reduction further enhances predictive performance. These findings imply that post-processing techniques, such as noise filtering and dimensionality reduction, are crucial for enhancing the utility of text indicators and can contribute to improving the accuracy of economic forecasts.

Keywords: big data, text data, sufficient dimension reduction, forecasting model

1. 서론

경제분석에 비정형 데이터를 활용하려는 시도가 크게 증가하고 있다. 그러나 통계 공표기관에서 작성하는 공식 통계와 달리 비정형 데이터는 노이즈를 포함하고 가공 방식에 따라 정보의 양이 달라질 수 있다는 우려가 제기된다. 이는 공식통계가 엄격한 품질 관리를 통해 기획되는 조사 방법론(experimental study)에 기반을 두는 반면, 비정형 데이터는 목적 없이 기록된 데이터에서 정보를 추출하는 관측 기법(observational study)에 기반을 두기 때문이다. 특히 경제적으로 관심이 높은 비정형 데이터는 음성, 텍스트, 이동(mobility) 정보 등으로, 특정되지 않은 방대한 양의 정보를 포함하는 빅데이터라는 점에서 정보를 가공하는 방식이 정보의 질에 영향을 미칠 것임은 자명하다.

최근의 연구 Seo (2023a, 2023b)와 Seo 등 (2024)에 의하면 뉴스기사, 증권사 리포트 등 텍스트 정보가 경제 분석에 매우 유용함을 알 수 있다. 그러나 해당 연구에서 추출한 텍스트 지표들은 변동성이 다소 높게

This work was supported by the Bank of Korea.

Jae Keun Yoo and Yujin Park are co-first authors. For Jae Keun Yoo and Yujin Park, this work was supported by Basic Science Research Program through the National Research Foundation of Korea (NRF) funded by the Korean Ministry of Education (RS2023-00240564 and RS-2023-00217022).

¹Corresponding author: Office of Economic Modeling and Policy Analysis, Bank of Korea, 39 Namdaemun-ro, Jung-gu, Seoul 04531, Korea. E-mail: bsseo@bok.or.kr

Published 30 April 2024 / journal homepage: <http://kjas.or.kr>

© 2024 The Korean Statistical Society. All rights reserved.

나타난다. 이는 텍스트 지표 작성을 위해 뉴스기사 등에 산재해 있는 정보를 추출하고 합성하는 과정에서 지표에 일부 노이즈가 포함되거나 정보가 효율적으로 가공되지 않았을 가능성이 있음을 시사한다. 따라서 본 연구는 노이즈 필터링 등의 기법과 차원축소 등의 방법을 이용하여 텍스트 지표의 정상화에 대해 검토하고 후처리 과정의 실증분석을 통해 텍스트 지표의 활용가능성을 높이고자 하였다. 특히 Hodrick과 Prescott (1997) 필터 등을 이용하여 텍스트 지표의 노이즈를 제거한 경우를 비교하고, 특정 조절 모수 값이 국내총생산(gross domestic product; GDP), 기업경기실사지수(business survey index; BSI) 등 경제지표 예측에 보다 유리한 결과를 제공하는지를 실증분석을 통해 검토하였다. 또한 텍스트는 방대한 양의 정보를 포함하므로 다양한 경제 지표를 산출할 수 있는 점을 고려하여, 해당 정보들의 차원 축소를 통해 정보를 합성하는 경우 경제 지표 예측에 보다 유리한 결과를 제공하는지를 충분차원축소(sufficient dimension reduction; SDR) 등을 통해 검증하였다. 또한 Seo (2023b)가 산출한 산업별 네트워크 데이터와 개별 산업 텍스트 지표의 정보량을 비교하고 텍스트에서 추출한 네트워크 데이터를 경제분석에 활용하는 방안을 함께 검토하였다.

텍스트 지표를 통계 모형에 활용하기 위해서는 텍스트에 포함된 광범위한 정보 중 목적 변수와 상관관계가 높은 주요 지표를 식별하는 과정이 중요하다. 이를 위해 본 연구에서는 경제 변수와 관련이 높은 정보는 유지하면서, 고차원의 자료를 저차원의 자료로 변환하는 차원축소 방법론을 주요하게 검토하였다. 특히 시계열 데이터의 특성을 감안하여 텍스트 정보가 거시 경제변수에 반영되어 나타나는 시차를 고려할 수 있도록 시점 정보를 고려한 차원축소 방법론을 제시하였다. 또한 노이즈 필터링 기법은 차원축소를 위한 전처리 과정으로 간주하고 노이즈 필터링과 차원축소 방법을 연결하여 분석하였으며, 경기 예측력을 기준으로 각각의 방법론들을 평가하였다.

본 연구는 최근 다양하게 연구되고 있는 텍스트 데이터의 필터링과 차원축소 등을 통하여 비정형 데이터의 활용가능성을 높이는 방안을 다각도에서 검토했다는 점에서 의의가 있다. 또한 본 연구에서 제시한 필터링 정보와 시차를 고려한 차원축소 등의 방법론은 텍스트 이외의 다양한 비정형 경제 데이터에도 비슷하게 적용될 수 있다는 점에서 향후 관련 연구에도 기여할 수 있을 것으로 기대된다.

본 논고는 다음과 같이 구성하였다. 2장에서 관련 선행연구와 연구배경에 대해 살펴보고 이어지는 3장에서 실증분석에 활용한 데이터를 소개하였다. 4장에서는 적용 방법론을 자세히 설명하였고, 5장에서 실증분석 결과를 제시하였다. 마지막으로 6장에서 본 연구의 시사점과 발전방향을 정리하였다.

2. 관련 선행연구 리뷰와 연구배경

2.1. 선행연구

텍스트 지표의 유용성에 대한 연구는 최근 활발히 진행되고 있다. Seo (2023a)는 비정형 뉴스 텍스트의 정량화를 통해 경제지표를 작성하기 위한 새로운 텍스트 마이닝 방법론을 제시하고 있다. 특히 다양한 분야의 경제뉴스 중에서 관심이 높은 생산, 물가, 고용, 주가, 주택 가격 등 15개 분야에 대해 정량화된 텍스트 지표를 제시하고 동적인자모형(dynamic factor model)과 합성곱 순환신경망(convolution recurrent neural network; CRNN)을 이용하여 텍스트 지표의 경기예측력을 평가하였다. Seo (2023a)의 분석결과에 따르면 텍스트 지표와 공식 통계를 함께 사용하여 경기 예측 모형을 구성할 경우 GDP의 평균 예측정확도가 향상되는 것으로 나타난다. 해당 선행연구에서 제시한 15개의 텍스트 지표는 월별 지표로 산출되었다.

또한 Seo (2023b)는 증권사 기업평가 보고서의 텍스트 자료를 활용하여 산업별, 지역별 업황 분석을 시도하였다. 해당 선행연구에서는 2019년도 이후 증권사 기업 평가 보고서 약 13만건을 빅데이터로 구축하고 이로부터 산업별 업황지수, 산업간 네트워크 지표 등 정량화된 텍스트 지표들을 산출하였다. 증권사 애널리스트 보고서의 텍스트 정보는 시점별 자료의 양, 시계열의 변동성, 계산 편의 등을 고려하여 분기별로 산출되었다. 증권사 애널리스트 보고서의 통계적 분석을 위해 정량 텍스트정보를 추출하는 자세한 전처리 과정과 변환 방법에 대해서는 Seo (2023b)의 3절을 살펴보기를 바란다. 해당 선행연구 분석결과에 따르면 업종별

텍스트 업황은 관련 코스피 산업별 지수 및 거시 경제지표와 비슷한 패턴을 보이고, 1분기 내외의 선행성을 갖는 것으로 나타난다. 또한 동 텍스트 지표들은 관련 코스피 산업별 지수에 대해서도 0.5~0.9의 높은 상관 관계를 보인다. 한글이 아닌 다른 언어의 텍스트 분석과 관련된 해외 연구사례는 Seo (2023b)의 1절에 자세히 소개되어 있다.

한편 텍스트 데이터의 노이즈 필터링 등 정상화에 대한 연구는 아직 시도된 바가 많지 않다 (Chen 등, 2021). 이는 경제 분석에 텍스트 데이터가 활용되기 시작한지가 오래되지 않은 점 때문으로 사료된다. 텍스트 데이터 이외의 다양한 금융경제 데이터에 대한 노이즈 필터링 연구를 살펴보면, 실용적 관점에서는 Hodrick과 Prescott (1997)의 방법론이 가장 널리 활용되고 있다. HP 필터는 경제 변수의 추세변동과 순환변동을 구분하기 위해 제시된 비모수적 방법으로 다음의 산식을 이용하여 산출한다:

$$\min_{\tau} \left(\sum_{t=1}^T (y_t - \tau_t)^2 + \lambda \sum_{t=2}^{T-1} [(\tau_{t+1} - \tau_t) - (\tau_t - \tau_{t-1})]^2 \right),$$

여기서 y_t 는 관측변수를 나타내며 y_t 는 추세 요소 τ_t 와 순환성 요소 c_t , 및 잔차 요소 ϵ_t 의 합으로 표시되는 일반적인 가법 모형에 의해 분해된다. 여기서 추세 요소는 시계열의 장기적인 움직임을 나타내며, 순환성 요소는 추세와 관련이 없는 정기적인 패턴을 그리고 잔차 요소는 추세나 순환성 요소로 설명할 수 없는 무작위한 변동을 나타낸다. HP필터에서 조절 모수값 λ (lambda)는 추세 요소의 변동성을 결정하는 초매개변수(hyperparameter)로 동 조절 모수값이 커질 수록 매끄러운(smooth) 추세 요소를 얻을 수 있다. Hodrick과 Prescott (1997)과 Ravn과 Uhli (2002)는 월자료에 대해 129,600의 λ 를 그리고 분기 자료에 대해 1,600의 λ 를 경험 법칙(heuristics)으로 제시하였다.

그러나 텍스트 지표의 변동성이 여타 금융 지표보다 높게 나타나며, 또한 본 연구의 목적이 순환성 요소를 찾기 보다는 노이즈 제거에 있다는 점에서 HP 필터를 이용할 경우 새로운 lambda 값을 탐색하는 과정이 필요할 것으로 판단된다. HP 필터 이외에 Kalman 필터, Hamilton (2018) 필터 등의 방법론도 경제지표의 필터링 방법으로 널리 활용되고 있는데, 추정과정이 HP 필터에 비해 복잡하고 컴퓨팅 시간이 오래 걸리므로 본 연구에서는 고려하지 않았다.

고차원 자료에서 정보를 가공하는 방식으로는 필터링 외에 차원축소의 방법도 고려할 수 있다. 차원 축소는 일반적인 회귀 문제에서 예측력을 높이는 데 이바지할 수 있는 것으로 알려져 있다. Yoo (2019)는 초고차원 자료의 대표인 마이크로 어레이 자료에 대한 생존분석에서 유전체 자료를 차원축소를 통하여 콕스-비례 모형을 적합할 때 위험율(hazard ratio) 대한 예측력을 높일 수 있음을 보여주고 있다. 총 7,399의 유전체 정보를 우선 주성분 분석을 통하여 40개로 일차적으로 차원 축소를 하였다. 훈련 자료에 대한 표본의 수가 160개 이기에 여전히 40개의 변수는 적지 않은 변수의 수이다. 이를 다시 fused sliced inverse regression을 적용하여 1차원으로 축소하여, 최종적으로는 7,399차원의 유전체 자료를 1차원의 자료로 축소하였다. 이렇게 최종 축소된 자료를 이용하여 생존 분석에서 가장 보편적으로 많이 사용되는 콕스-비례 모형에 적합을 하여 실제 예측력이 기존보다 높아 졌음을 검증 자료를 통하여 보여주었다. 차원 축소가 자료분석을 보다 정확하게 해 줄 수 있음은 이러한 의생명 분야뿐만 아니라 화학 분야 (Lee 등, 2019), 식품 분야 (Kim 등, 2017) 등 다양한 분야에서 실증되고 있다. 이러한 충분차원축소 방법론을 본 연구에서는 비정형 경제 데이터인 텍스트 데이터에 적용하였다. 특히 본 연구에서는 sliced inverse regression (Li, 1991), directional regression (Li와 Wang, 2007)과 seeded method (Cook 등, 2007)를 고려하고자 한다.

2.2. 연구배경

COVID-19, 러우전쟁 등 외생적인 경제 이벤트는 소비, 투자 등 경제주체의 경제활동에 영향을 미친다. 그러나 과거 데이터를 기반으로 한 실증 분석을 통해서도 경제 이벤트의 정확한 효과를 파악하는 것이 쉽지 않다.

이는 경제적 영향이 큰 이벤트일 수록 과거에 유사한 사례를 찾는 것이 어렵기 때문이다. 따라서 경제가 급변하는 시기에는 뉴스 등 정성적 평가를 반영하는 비정형 데이터의 가치가 크게 증가한다. 뉴스와 같은 텍스트 자료의 경우 실시간으로 다양한 소스를 통해 정보가 생성되며, 전문가 판단 등을 반영하므로 정보의 양과 질 면에서 경제 이벤트를 설명하는 풍부한 정보를 제공한다고 할 수 있다.

국내에서 경기 판단을 위한 경기선행지수, BSI 및 GDP 등의 예측을 위해 비정형 텍스트 자료를 이용한 분석은 아직 많이 시도되지 않고 있다. 그러나 비정형 텍스트 자료가 가지는 정보의 다양성과 방대함 그리고 속보성을 고려할 때 비정형 데이터를 통계 모형에 반영하여 분석할 필요성은 충분하다.

이러한 가능성을 배경으로 본 연구에서는 경기선행지수, BSI, GDP 등의 목표 변수에 대해 최근 한국은행에서 개발되어 사용되고 있는 뉴스 기반 경제 부문별 텍스트 자료, 애널리스트 리포트 기반 텍스트 자료 그리고 애널리스트 리포트 기반 네트워크 자료 등을 추가한 경기 예측 모형을 개발 및 평가하는 것을 그 목표로 두고자 한다. 구체적으로 비정형 텍스트 자료의 변동성 안정을 위한 자료변환, 차원 축소 등 텍스트 지표의 후처리 과정이 실제 분석에서 예측 정확도를 높일 수 있는지 실증분석을 통해 검증하고 이를 통해 비정형 텍스트 지표의 활용 가능성을 제고하고자 한다.

3. 자료 소개

3.1. 텍스트 지표

3.1.1. 뉴스 기반 경제 부문별 텍스트 자료

2005년 이후 뉴스 기반 경제 부문별 텍스트 지표(theme frequency indices; TFI)로 Table A.1과 같이 주요 거시 변수 및 일부 산업 관련 미시 변수를 선정하여 15개 부문으로 구성되어 있다. 본 연구에서 사용되는 TFI 자료는 Seo (2023a)에 기반한 자료이다. TFI 자료의 경우 월별 자료이고, 분기별 분석을 위해서는 1~3월을 1분기, 4~6월을 2분기, 7~9월을 3분기, 10~12월을 4분기로 합성하여 월 자료의 평균치를 각 분기 값으로 사용하였다.

3.1.2. 애널리스트 리포트 기반 텍스트 업황 자료

2019년 이후 애널리스트 리포트 기반 산업별 텍스트 지표의 업황 분석(market analysis; MA) 자료로 Table A.2와 같이 41개 부문으로 구성되어 있다. MA 자료는 앞서 언급된 선행 연구인 Seo (2023b)에 기반한 자료이다. MA 자료의 경우는 월별, 분기별 자료로 구분되어 있다.

3.1.3. 네트워크 지표

2019년 이후 애널리스트 리포트 기반 자료에서 추정된 산업별 유사도 네트워크 자료로 Table A.3과 같이 39개 부문으로 구성되어 있고, 이는 MA 자료의 연번 1~39번과 동일하다. MA 자료 연번의 40~41번은 결측치가 과다하게 나타나 동 자료 작성에서는 제외되었다.

3.2. 예측 목표 변수

월별 자료에 대한 예측 목표 변수는 선행지수 순환 변동치, BSI 전산업 매출실적, BSI 전산업 매출전망이다. 분기별 자료의 예측 목표 변수는 GDP 실질 SA전기비와 GDP 실질 원계열 전년 동기비이다.

4. 자료의 차원 축소

본 연구에서 사용되는 변수는 TFI 자료와 MA 자료인데 두 자료의 변수의 수는 각각 15개와 41개이다. 2005년부터 월별로 만들어진 TFI 자료의 경우 15개의 변수의 수는 그렇게 높은 차원이라고 할 수는 없지만, 2019년

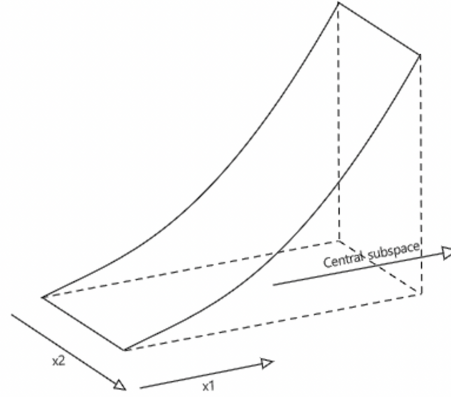


Figure 1: Central subspace; X_1 and X_2 are two dimensional covariates.

부터 분기별로 작성된 MA 자료의 경우 자료의 수에 비해 41개의 변수는 비교적 많은 편이다. 이러한 고차원 분석에서 변수의 차원 축소는 모형을 보다 안정적으로 하기 위해 중요한 역할을 할 수 있음이 이미 다양한 형태의 고차원 자료분석에서 제시되고 있다. 이를 위해 본 연구에서는 충분차원축소(sufficient dimension reduction)의 접근법을 사용하고자 한다.

4.1. 충분차원축소

충분차원축소(sufficient dimension reduction)란 회귀분석에서 고차원 데이터의 차원을 줄이기 위해 사용되는 통계적 방법이다. 충분차원축소는 주로 회귀분석에서 사용되는 고차원의 설명변수의 차원을 축소하는 것이 그 주요 목적이지만, 반응 변수의 차원이 고차원인 경우 반응 변수의 차원도 축소도 그 목적에 포함된다. 하지만 본 연구에서는 일 차원의 반응 변수를 고려하므로 반응 변수의 차원 축소에 대해서는 언급하지 않을 것이다. 이하 내용에서의 차원 축소는 회귀분석에서 설명 변수의 차원 축소에 대한 것임을 먼저 밝혀 두고자 한다.

여러 가지 특징이 있는 고차원의 데이터를 사용하는 경우 관심 있는 반응 변수를 예측하는 데 관련성이 높은 특징을 식별해 내는 것이 어려움에 직면할 수 있다. 충분차원축소의 목표는 회귀분석을 보다 정확히 하기 위해 관련된 설명 변수 정보는 유지하면서, 원 설명 변수의 저차원의 선형 결합 형태로 차원을 축소하는 방법론이다. 이를 보다 구체적으로 설명하면 원래 p 차원의 설명 변수 $\mathbf{X} \in \mathbb{R}^p$ 가 주어졌을 때 이를 회귀의 목적인 조건부 분포 $\mathbf{Y}|\mathbf{X}$ 에 대해 정보의 손실 없이 원래의 설명 변수를 저차원의 선형변환인 $\mathbf{C}^T\mathbf{X}$ 로 대체하는 것을 의미한다. 여기서 $\mathbf{C}$ 는 $p \times q$ 차원의 행렬이고, 이는 다음과 같이 표현할 수 있다:

$$\mathbf{Y} \perp\!\!\!\perp \mathbf{X} | \mathbf{C}^T\mathbf{X},$$

여기서 $\perp\!\!\!\perp$ 는 통계적 독립을 뜻한다. 만약 $q \leq p$ 성립한다면, 원래 p 차원의 설명 변수 $\mathbf{X}$ 는 회귀에 대한 정보의 손실없이 q 차원의 $\mathbf{C}^T\mathbf{X}$ 로 대체될 수 있다. 위의 식을 만족하는 행렬 $\mathbf{C}$ 의 열에 의해 생성되는 공간을 차원 축소 부분 공간(dimension reduction subspace)이라고 한다.

하나의 회귀분석 문제에서 위의 식을 만족하는 $\mathbf{C}$ 는 여러 개 존재할 수 있다. 그렇다면 차원 축소의 목적 상 최소 부분 공간을 찾는 것이 바람직하다. 여러 개의 차원 축소 부분 공간이 있다는 전제 하, 최소의 차원 축소 부분 공간은 이들의 교집합을 통해 만들어 질 수 있다. 즉, 가능한 모든 차원 축소 부분 공간들의 교집합이 차원 축소 부분 공간이 된다면, 당연히 최소 공간일 뿐만 아니라 유일성도 확보될 수 있다. 이 교집합의 차원

축소 부분 공간을 중심부분공간(central subspace)이라고 하며, $S_{Y|X}$ 로 정의한다. 충분차원축소의 주요 목적은 당연히 이 중심부분공간을 추정하는데 있다. 중심부분공간의 차원을 d 라고 표시할 것이며, 정치 기저 행렬을 $\eta \in \mathbb{R}^{p \times d}$ 으로 나타낼 것이다. 또한 저 차원의 선형 결합 형태인 설명변수 $\eta^T X$ 를 충분축소변수(sufficient predictor)라고 할 것이다.

방법론적으로 충분차원축소의 핵심 아이디어는 고차원 특징 공간의 선형 또는 비선형 투영을 저차원 부분 공간에 찾는 것으로 투영은 관련 없는 정보를 제거하면서 가능한 많은 데이터 구조를 보존해야 한다. 이는 Figure 1 (Li, 2018) 과 같이 반응 변수를 설명하기에 충분한 공간의 방향 집합을 찾아냄으로써 달성할 수 있다. 방향 집합을 식별함으로써 전체 데이터가 아닌 소수의 특징 데이터만을 이용하여, 통계 모델의 정확성과 해석 가능성을 향상시킬 수 있으며, 이는 시각화 측면에도 도움이 될 수 있다. 입력 데이터가 고차원적이고 관측 횟수가 제한적인 상황에서 특히 유용한데, 이는 과적합 위험을 줄이고 통계 추론의 효율성을 향상시키는 데 도움이 될 수 있기 때문이다. 이러한 특징들로 인해 충분차원축소(SDR) 방법은 많은 분야에서 적용되고 있으며, 특히 고차원 데이터가 주로 쓰이는 유전학, 신경과학, 금융 등의 분야에서 사용되고 있다.

충분 차원 축소와 중심부분공간의 정의에서 회귀 모형에 대한 특정한 가정을 하지 않고 있다. 즉 충분 차원 축소는 방법론적으로 모형을 기반한 모수적 방법론보다는 중심부분공간을 추정하기 위한 최소의 가정들만으로 사용될 수 있는 비모수 혹은 반모수적 방법론들이 보편적이다. 보통 중심부분공간을 추정할 때 $Y|X$ 에서 직접 추정하기 보다는 $Y|Z$ 에서 회귀를 고려하여 추정한다. 행렬 Σ 를 X 의 공분산 행렬이라고 정의하자. 그러면 Z 은 다음과 같이 정의 된다:

$$Z = \Sigma^{-\frac{1}{2}} (X - E(X)).$$

즉 Z 는 원래의 설명 변수를 평균이 0이고 공분산 행렬이 자기 행렬(identity matrix)이 되도록 변환한 설명 변수이다. 그렇다면 $Y|X$ 의 중심 부분 공간과 $Y|Z$ 의 중심 부분 공간 사이에는 다음이 성립한다:

$$S_{Y|X} = \Sigma^{-\frac{1}{2}} S_{Y|Z}.$$

표준화한 설명 변수 Z 를 고려하는 주요 이유는 추정에 있어 발생하는 계산 오차를 줄여 보다 정확하게 중심 부분 공간을 추정하는데 있다. 이후 $S_{Y|Z}$ 의 중심부분공간에 대한 정치 기저 행렬을 η_z 라고 정의한다. 이후 소개되는 충분차원축소 방법들은 모두 방법론적 소개를 $Y|Z$ 회귀를 고려하여 하고자 한다. 지난 20년 넘게 중심 부분 공간을 추정하는 다양한 방법론들이 개발되었지만, 본 연구에서는 다음의 대표적인 3가지 방법을 고려할 것이다.

- (1) Sliced Inverse Regression (SIR) (Li, 1991).
- (2) Directional Regression (DR) (Li와 Wang, 2007).
- (3) Seeded Method (seed) (Cook 등, 2007).

위 세 가지 방법들을 다음 세 항에 걸쳐 소개하고자 한다.

4.1.1. Sliced inverse regression

Sliced inverse regression (SIR) 방법은 충분 차원 축소 방법 중 가장 오래된 역사를 지닌 방법 중 하나이다. 1991년 Li에 의하여 처음으로 소개되었으며, 그 이후로 다양한 응용 분야에서 사용되고 있다. 널리 알려진 차원 축소 방법 중 하나인 주성분분석(principal component analysis; PCA) 방법과의 차이점은 PCA 방법은 차원 축소를 진행할 때 반응 변수가 축소과정에서 사용되지 않지만, SIR 방법은 반응 변수가 사용된다. 이는 SIR가 반응 변수와 관련 있는 설명 변수를 식별하는데 더 유용할 수 있고, PCA보다 회귀분석에서 차원 축소를 보다 효과적으로 수행할 수 있음을 암시하고 있다.

SIR가 중심 부분 공간을 추정하기 위해서는 다음의 선형성 조건이 만족되어야 한다:

$$E(\mathbf{Z} | \boldsymbol{\eta}_\epsilon^T \mathbf{Z}) \text{는 } \boldsymbol{\eta}_\epsilon^T \mathbf{Z} \text{에 대하여 선형이다.}$$

위의 조건은 설명 변수 $\mathbf{Z}$ 에 대해서만 적용되는 조건이다. 위의 선형성 조건은 중심부분공간을 생성하는 기저 $\boldsymbol{\eta}_\epsilon$ 에 대해서만 성립하면 된다. 이외의 다른 충분 차원 축소 공간의 기저에 대해 성립할 필요는 없다. 따라서 위의 조건은 매우 강한 조건으로 보일 수 있다.

실제로 회귀분석에서의 분포적 혹은 이외 필요한 다양한 조건들은 조건부 분포인 $\mathbf{Y}|\mathbf{X}$ 에 가정된다. 또한 일반적으로 회귀분석에서는 설명 변수에 대해서는 별다른 조건을 부여하지 않는다. 하지만 보다 나은 적합을 위해 대부분의 경우 설명 변수가 다변량 정규분포를 따르도록 Box-Cox 변환 방법을 사용한다. 왜냐 하면 그러한 경우 다변량 정규분포 이론에 의해 $\mathbf{Y}|\mathbf{X}$ 는 정규분포를 따르게 되고, 선형 회귀의 적합이 자료를 잘 설명할 수 있기 때문이다.

위의 선형성 조건은 설명 변수 $\mathbf{Z}$ 의 분포가 타원대칭분포(elliptically contoured distribution)를 따른다면 모든 선형변환 $\boldsymbol{\eta}_\epsilon^T \mathbf{Z}$ 에 대해 성립한다. 이러한 타원대칭분포 중 하나가 다변량 정규분포이고, 위에서 논의된 바와 같이 설명 변수의 다변량 정규분포 가정은 실제 회귀분석에서 보편적으로 가정되고 있다. 또한 Hall과 Li (1993)에 따르면 비록 $\mathbf{Z}$ 가 타원대칭분포를 따르지 않을지라도 설명 변수의 차원이 높아지면 이 선형성 조건은 만족된다는 것을 증명하였다. 만약 선형성 조건을 만족시키고자 한다면 설명 변수가 다변량 정규분포를 따르도록 변수 변환을 하는 것이 일반적이다. 이는 선형성 조건이 조건 자체가 분포를 가정하는 것처럼 강하지 않으며, 조건이 만족하지 않더라도 현실적인 대안이 있다는 점을 강조하고자 한다. Li (1991)에 따르면 위의 선형성조건이 만족되면, 다음의 식이 성립한다:

$$\Sigma^{-\frac{1}{2}} E(\mathbf{Z} | \mathbf{Y}) \subseteq S_{\mathbf{Y}|\mathbf{Z}}.$$

중심부분공간 $S_{\mathbf{Y}|\mathbf{Z}}$ 의 추정은 본질적으로 역회귀인 $E(\mathbf{Z}|\mathbf{Y})$ 를 추정함으로써 가능하다. 방법론에서 분포적 가정을 하지 않았기 때문에 이 추정은 비모수적일 수 밖에 없다. $E(\mathbf{Z}|\mathbf{Y})$ 는 만약 반응 변수가 이산형이라면 직관적이다. 반응 변수가 이산형일 경우 $E(\mathbf{Z}|\mathbf{Y})$ 는 단순히 반응 변수의 범주에 해당되는 $\mathbf{Z}$ 의 표본 평균으로 추정될 수 있다. 반응 변수가 연속형인 경우 해당 반응변수를 이산형으로 범주화한다면 쉽게 $E(\mathbf{Z}|\mathbf{Y})$ 를 추정할 수 있을 것이다. Li (1991)에서는 이러한 연속형 반응변수의 그룹화를 슬라이싱(slicing)이라고 표현하였다. 슬라이싱을 할 때 자료가 겹치지 않도록 고르게 나누는 것은 중심부분공간의 정확한 추정을 위해 중요하다.

SIR는 일반적으로 선형적추세를 보이는 회귀분석 문제에서 잘 작동하는 것으로 알려져 있다. 해당 방법론 적용 단계를 알고리즘으로 정리하면 Algorithm 1과 같다.

요약하자면, SIR는 설명 변수와 반응 변수 사이의 관계에 대한 가장 많은 정보를 포착하는 역회귀 함수의 공간에서 가장 중요한 방향을 식별하는 데 좋은 성능을 나타내는 방법이다. 설명 변수를 이러한 방향으로 투영함으로써, 반응 변수에 대한 대부분의 정보를 보존하는 설명 변수의 저차원 표현식을 얻을 수 있다. SIR는 계산의 효율성으로 인해 변수가 아주 많은 고차원 데이터를 처리할 수 있다는 장점이 있으며, 데이터 분포에 대한 가정에 의존하지 않으므로 다양한 데이터에 적용이 가능한 방법이다. 하지만 SIR의 한계는 슬라이싱 개수에 민감하다는 점이며, 중심 부분 공간이 정확하게 식별되도록 충분한 수의 방향을 선택하는 것이 중요하다.

4.1.2. Directional regression

SIR는 보편적으로 많이 사용되는 충분차원축소 방법이지만 위에서 언급된 방법론적 한계가 명백히 존재한다. 특히 자료가 반응 변수와 설명 변수의 관계가 다음과 같은 $y = x^2$ 과 같은 대칭 관계를 가질 때 SIR는 이러한 관계를 추정할 수 없다. 이 대칭 관계는 second-order polynomial 회귀로 실제 자료분석에 자주 사용되는 선형 회귀 모형 중 하나이다.

Algorithm 1 : Sliced inverse regression

1. 설명변수의 표본 평균과 표본공분산을 계산한다.

$$\hat{\mu} = E_n(\mathbf{X}), \quad \hat{\Sigma} = \text{cov}_n(\mathbf{X}).$$

2. 설명 변수를 표준화한다.

$$\mathbf{Z}_i = \hat{\Sigma}^{-\frac{1}{2}} (\mathbf{X}_i - \hat{\mu}), \quad i = 1, \dots, n.$$

3. 반응변수가 연속형인 경우 $\mathbf{Y}$ 를 h 개의 분할구간 $j_1, \dots, j_h$ 로 나누고, 각 분할구간에서의 표준화된 설명변수의 평균을 구한다.

$$E_n(\mathbf{Z} | \mathbf{Y} \in J_l) = \frac{E_n[\mathbf{Z}I | \mathbf{Y} \in J_l]}{E_n[I | \mathbf{Y} \in J_l]}, \quad l = 1, \dots, h.$$

4. 공분산 행렬을 계산한다.

$$\hat{\Lambda} = \sum_{l=1}^h E_n[I(\mathbf{Y} \in J_l)] E_n(\mathbf{Z} | \mathbf{Y} \in J_l) E_n(\mathbf{Z}^T | \mathbf{Y} \in J_l).$$

5. $\hat{\Lambda}$ 의 고유벡터 중 가장 큰 고유치 d 개의 고유벡터를 $\hat{v}_1, \dots, \hat{v}_d$ 이라고 정의하면, $\hat{\eta}_z = (\hat{v}_1, \dots, \hat{v}_d)$ 이고, $\hat{\eta} = \hat{\Sigma}^{-1/2}(\hat{v}_1, \dots, \hat{v}_d)$ 이 된다.

SIR가 가지는 또 다른 방법론적 한계는 중심부분공간을 완전히 추정하지 못하는 것이다. 이를 위해 Li 등 (2005)은 contour regression을 제안하였다. Contour regression은 완전하게 중심부분공간을 추정할 수 있는 장점이 있지만, 모든 contour의 방향을 계산해야하는 복잡성으로 인해 상당한 시간이 소요된다는 결정적 단점이 있다. Contour regression의 이러한 계산적 복잡성을 보완하기 위하여 개발된 directional regression (DR) (Li와 Wang, 2007)을 본 연구에서는 고려하고자 한다. DR은 역 회귀의 두 조건부 적률을 사용하여 중심부분공간을 추정하는 방법으로써 반응 변수와 설명 변수 간의 공분산을 최대화하는 방향 또는 축의 집합을 찾는 것을 기반으로 한다. DR은 역회귀의 2차적률도 고려함으로써 contour regression이 가지는 계산적 복잡함을 해결하였으며, SIR가 가지는 비선형 회귀에서의 문제점도 완화할 수 있다. 또한 DR은 SIR가 가지는 slicing의 민감함에 강건하다는 장점이 있다. DR은 SIR과 같이 예측 모델의 정확도를 향상시키기 위한 전처리 단계로 사용될 수 있으며, DR은 다른 차원 축소 기법보다 이상치 및 노이즈에 덜 민감하고, 계산 효율성이 높은 방법으로 알려져 있다.

DR을 사용하기 위해서는 SIR에서 필요한 선형성 조건 외에 다음의 등분산 조건을 만족해야 한다:

$$\text{cov}(\mathbf{Z} | \eta_z^T \mathbf{Z}) \text{는 } \eta_z^T \mathbf{Z} \text{에 대해 일정하다.}$$

등분산 조건은 선형성 조건과 마찬가지로 설명 변수에만 해당되는 조건이다. 하지만 선형성 조건과는 다르게 타원대칭분포에서 항상 성립하지 않지만, 다변량 정규분포에서는 만족한다.

Li와 Wang (2007)에 따르면 선형성조건과 등분산 조건이 만족되면, 다음의 식이 성립하고, 아래의 식을 규명하는 것이 방법론적으로 DR이 가지는 가장 핵심적인 부분이다:

$$(2I_p - E[(\mathbf{Z} - \bar{\mathbf{Z}})(\mathbf{Y} - \bar{\mathbf{Y}})])^2 \subseteq S_{\mathbf{Y}|\mathbf{Z}},$$

여기서 $\bar{\mathbf{Z}}$ 와 $\bar{\mathbf{Y}}$ 는 각각 $\mathbf{Z}$ 와 $\mathbf{Y}$ 의 독립적 복사 변수들이다.

자료를 이용한 $(2I_p - E[(\mathbf{Z} - \bar{\mathbf{Z}})(\mathbf{Y} - \bar{\mathbf{Y}})])^2$ 의 추정은 SIR과 같이 먼저 연속형 반응 변수를 슬라이싱을 통하여

Algorithm 2 : Directional regression

1. 설명변수의 표본 평균과 표본공분산을 계산한다.

$$\hat{\mu} = E_n(\mathbf{X}), \quad \hat{\Sigma} = \text{cov}_n(\mathbf{X}).$$

2. 설명 변수를 표준화한다.

$$\mathbf{Z}_i = \hat{\Sigma}^{-\frac{1}{2}} (\mathbf{X}_i - \hat{\mu}), \quad i = 1, \dots, n.$$

3. 반응변수가 연속형인 경우 $\mathbf{Y}$ 를 h 개의 분할구간 $j_1, \dots, j_h$ 로 나누고, 각 분할구간에서의 표준화된 설명변수의 평균을 구한다.

$$\mathbf{M}_{1l} = E_n(\mathbf{Z} | \mathbf{Y} \in J_l); \mathbf{M}_{2l} = E_n(\mathbf{Z}\mathbf{Z}^T | \mathbf{Y} \in J_l), \quad l = 1, \dots, h.$$

4. 이를 이용하여 다음의 3개의 행렬을 계산한다;

$$\begin{aligned} \Lambda_1 &= \sum_{l=1}^h \hat{p}_l \mathbf{M}_{1l}^2; \\ \Lambda_2 &= \left(\sum_{l=1}^h \hat{p}_l \mathbf{M}_{1l} \mathbf{M}_{1l}^T \right)^2; \\ \Lambda_3 &= \left(\sum_{l=1}^h \hat{p}_l \mathbf{M}_{1l}^T \mathbf{M}_{1l} \right) \left(\sum_{l=1}^h \hat{p}_l \mathbf{M}_{1l} \mathbf{M}_{1l}^T \right), \end{aligned}$$

여기서 $\hat{p}_h$ 는 h 번째 수준에서 관측수의 비율을 의미한다.

5. 4번째에서 계산된 행렬을 이용하여 다음의 행렬을 계산한다.

$$\Lambda_{DR} = 2\Lambda_1 + 2\Lambda_2 + 2\Lambda_3 - 2I_p.$$

6. Λ_{DR} 의 고유벡터 중 가장 큰 고유치 d 개에 대한 고유벡터를 $\hat{v}_1, \dots, \hat{v}_d$ 이라고 정의하면, $\hat{\eta}_z = (\hat{v}_1, \dots, \hat{v}_d)$ 이고, $\hat{\eta} = \hat{\Sigma}^{-1/2}(\hat{v}_1, \dots, \hat{v}_d)$ 이 된다.

범주화 한 후 각 범주별 설명변수의 평균과 공분산 행렬을 계산함으로써 가능해진다. DR을 이용한 중심부분 공간의 추정 Algorithm 2와 같이 정리된다.

4.1.3. Seeded method

앞에서 설명된 두 가지 방법뿐만 아니라 대부분의 충분차원축소 방법들은 방법론적으로 설명 변수에 대한 표본 공분산 행렬의 역행렬 계산이 필수적이다. 하지만 자료의 수보다 변수의 개수가 많은($n \leq p$) 경우에는 표본 공분산 행렬의 역행렬이 존재하지 않는다. 이것은 충분차원축소의 목적을 고려한다면 아이러니한 것이다. 이를 극복하기 위하여 Cook 등 (2007)은 역행렬의 계산을 요구하지 않는 방법론을 제안하였고, 이 방법론을 seeded method라고 하고자 한다.

Seeded method은 방법론 이름 그대로 아래의 관계를 만족하는 seed행렬이라고 불리는 행렬 $v \in \mathbb{R}^{p \times d}$ 을 먼저 찾아야 한다.

$$\Sigma^{-1}S(v) \subseteq S_{Y|X} \Leftrightarrow S(v) \subseteq \Sigma S_{Y|X},$$

여기서 $S(v)$ 는 v 의 열에 의해서 생성되는 공간을 의미한다.

앞에서 언급한 선형성 조건이 만족하는 상황에서 seed행렬 v 로 사용될 수 있는 행렬의 후보는 다음과 같다:

$$\begin{aligned} E(\mathbf{X}|\mathbf{Y}) : \Sigma^{-1}\{E(\mathbf{X}|\mathbf{Y}) - E(\mathbf{X})\} \in S_{\mathbf{Y}|\mathbf{X}} &\Leftrightarrow S\{E(\mathbf{X}|\mathbf{Y}) - E(\mathbf{X})\} \subseteq \Sigma S_{\mathbf{Y}|\mathbf{X}}; \\ \text{cov}(\mathbf{X}, \mathbf{Y}) : \Sigma^{-1}\text{cov}(\mathbf{X}, \mathbf{Y}) \in S_{\mathbf{Y}|\mathbf{X}} &\Leftrightarrow S\{\text{cov}(\mathbf{X}|\mathbf{Y})\} \subseteq \Sigma S_{\mathbf{Y}|\mathbf{X}}. \end{aligned}$$

위의 seed행렬 후보 중 $\Sigma^{-1}\{E(\mathbf{X}|\mathbf{Y}) - E(\mathbf{X})\}$ 는 SIR가 중심부분공간을 추정하기 위해 사용되는 행렬이며, $\Sigma^{-1}\text{cov}(\mathbf{X}, \mathbf{Y})$ 는 최소제곱 추정량이다. 본 연구에서는 $\text{cov}(\mathbf{X}, \mathbf{Y})$ 를 seed행렬로 사용한다. 그리고 중심부분공간을 포함하는 $\mathbf{M}_{\mathbf{Y}|\mathbf{X}}$ 이 존재한다고 가정하자. 그렇다면 다음의 관계가 성립한다:

$$\Sigma^{-1}S(v) \subseteq S_{\mathbf{Y}|\mathbf{X}} \subseteq \mathbf{M}_{\mathbf{Y}|\mathbf{X}} \Leftrightarrow S(v) \subseteq \Sigma \mathbf{M}_{\mathbf{Y}|\mathbf{X}}.$$

Seeded method은 직접적으로 $S_{\mathbf{Y}|\mathbf{X}}$ 을 추정하는 것이 아니라 seed행렬을 이용하여 설명 변수의 역행렬을 계산하지 않고 $\mathbf{M}_{\mathbf{Y}|\mathbf{X}}$ 의 기저를 추정하는데 있다. 그러면 그 기저에 대한 투영 행렬을 일반적인 내적 공간이 아닌 $\langle a, b \rangle_{\Sigma}$ 로 정의하는 Σ -내적 공간에서 정의되는 투영 행렬을 만들고 이를 이용하여 중심부분공간을 추정한다.

행렬 R 을 $\mathbf{M}_{\mathbf{Y}|\mathbf{X}}$ 의 기저 행렬이라고 하자. $\langle a, b \rangle_{\Sigma}$ 내적에 대해 $\mathbf{M}_{\mathbf{Y}|\mathbf{X}}$ 에 투영하는 직교 투영 연산자 $P_{\mathbf{M}_{\mathbf{Y}|\mathbf{X}}(\Sigma)} = R(R^T \Sigma R)^{-1} R^T \Sigma$ 이다. 이 직교 투영 연산자의 계산에 역행렬이 필요하지만 $p \times p$ 행렬이 아닌 $q \times q$ 행렬인 $R^T \Sigma R$ 의 역행렬이다. 만약 q 가 표본 수 n 보다 작은 경우 역행렬 계산이 가능해진다. 만약 $\Sigma^{-1}v = S_{\mathbf{Y}|\mathbf{X}}$ 라면 다음의 관계가 성립된다:

$$S_{\mathbf{Y}|\mathbf{X}} = P_{\mathbf{M}_{\mathbf{Y}|\mathbf{X}}(\Sigma)} S_{\mathbf{Y}|\mathbf{X}} \Sigma^{-1} S(v) = S \left(R \left(R^T \Sigma R \right)^{-1} R^T \Sigma \Sigma^{-1} v \right) = S \left(R \left(R^T \Sigma R \right)^{-1} R^T v \right).$$

따라서 $R(R^T \Sigma R)^{-1} R^T v$ 는 중심부분공간을 생성하는 기저 행렬이 되고, 이를 통하여 중심부분공간을 추정할 수 있다. $R^T \Sigma R$ 의 역행렬은 일반적으로 무어-펜로즈(Moore-Penrose inverse) 역행렬로 계산이 가능하다.

위에서 언급된 행렬 R 은 자료를 통하여 추정되어야 하며, 실제로 seed행렬과 설명 변수의 표본 공분산 행렬을 이용하여 추정한다. 행렬 R 을 추정할 때 R 의 열공간이 $S_{\mathbf{Y}|\mathbf{X}}$ 를 포함할 만큼 충분히 커야 하며 동시에 자료를 통해 계산이 용이할 수 있도록 합리적 작아야 한다. 이를 위해 Cook 등 (2007)은 행렬 R 의 추정에 대해 v 를 Σ 에 반복적 투영하는 방법을 제시하고 있다.

$$R_u \equiv (v, \Sigma v, \dots, \Sigma^{u-1} v), \quad u = 1, 2, \dots, u^*,$$

여기서 u^* 는 투영종료지수(termination index of projections)라고 한다. 모든 $u \geq 2$ 에 대해 $S(R_{u-1}) \subseteq S(R_u)$ 관계가 성립한다. 그러므로 적절한 $S(R_u) = S_{\mathbf{Y}|\mathbf{X}}$ 를 보장하는 가장 작은 u 를 찾는 것이 중요하다. 이에 대해서는 R_u 의 정보가 더 이상 증가하지 않는 점을 찾아 정한다. 관련한 자세한 내용은 Um 등 (2018)을 참고하길 바란다. 해당 방법론 적용 단계를 알고리즘으로 정리하면 Algorithm 3과 같다.

5. 예측 모형 개발 및 평가

5.1. 예측 모형 개발 절차

뉴스 텍스트 지표는 15개 부문, 애널리스트 텍스트 지표는 41개 부문에 대한 지표로써 본 보고서에서 다루는 고차원자료라고 할 수 있다. 이에 4장에서 소개된 충분차원축소 방법들을 적용하여 자료의 차원을 축소하였다. 충분차원축소 방법들은 결측치가 없는 완전자료에서 적용이 가능하기 때문에 우선 결측치가 많은 MA의

Algorithm 3 : Seeded method

1. v 로 사용될 수 있는 후보 행렬 중 하나의 커널 행렬을 선택한다. 본 연구에서는 $\text{cov}(\mathbf{X}, \mathbf{Y})$ 를 seed행렬로 사용한다.

2. 설명 변수의 표본 공분산 행렬 $\hat{\Sigma}$ 와 seed 행렬 $\hat{v}$ 를 계산한다.

3. 적당히 큰 값의 u 에 대해 다음의 $\hat{R}_u$ 를 계산한다.

$$\hat{R}_u = (\hat{v}, \hat{\Sigma}\hat{v}, \dots, \hat{\Sigma}^{u-1}\hat{v}).$$

4. 계산된 $\hat{R}_u$ 를 통해 u^* 를 정하고, 다음의 $\hat{R}_{u^*}$ 를 계산한다.

$$\hat{R}_{u^*} = (\hat{v}, \hat{\Sigma}\hat{v}, \dots, \hat{\Sigma}^{u^*-1}\hat{v}).$$

5. $\hat{R}_{u^*}(\hat{R}_{u^*}^T \hat{\Sigma} \hat{R}_{u^*})^{-1} \hat{R}_{u^*}^T \hat{v}$ 를 계산하여 중심부분공간의 기저를 추정한다.

자료는 차원축소 전에 선형보간법을 이용하여 완전 자료로 만들었다. 또한 본 연구의 중요한 목적 중의 하나가 텍스트 지표의 변동성 안정화 및 노이즈 제거가 예측모형 개발에 효과가 있는지를 검토하는 것이다. 이에 다양한 조절 모수(lambda, λ) 값을 지정하여 HP filter를 TFI와 MA자료에 적용하였다. 이때 조절 모수 값을 해석이 쉽도록 자료의 주기(frequency, ρ) 지표로 치환하여 사용하였다. 주기와 조절 모수 사이에는 다음의 식이 성립한다 (Hodrick과 Prescott, 1997):

$$\lambda = 6.25\rho^4.$$

월별 자료에서 사용된 주기 값은 0.5, 1, 1.5와 2이고, 분기별 자료에서는 0.25, 0.5, 0.75와 1이 사용되었다. 이렇게 결측치에 대해 선형보간 후 HP filtering을 한 TFI와 MA에 대해 각각 앞 장에서 소개된 충분차원축소 방법을 적용하였다.

월별 예측 목표 변수인 선행지수 순환 변동치, BSI 전산업 매출실적, BSI 전산업 매출전망에 대해서 TFI만 고려할 경우, 2005년부터의 자료가 있어 적지않은 표본 수이기 때문에 슬라이스의 수를 10으로 하여 SIR과 DR를 적용하여, 최대 2차원으로 축소한다. 이후 모형화에서는 1차원으로 축소된 자료, 그리고 2차원으로 축소된 자료가 따로 사용된다. 그리고 MA만 그리고 TFI와 MA를 모두 고려한 경우에는 2019년부터의 월별 자료이기 때문에 표본의 수가 변수의 수보다 상대적으로 많이 크지 않기 때문에 seeded method를 적용하였고, 이때 사용된 seed행렬은 $\text{cov}(\mathbf{X}, \mathbf{Y})$ 이다. 따라서 이 때 TFI와 MA는 각각 1차원으로 축소가 된다.

분기별 예측 목표 변수인 GDP 실질 SA전기비와 GDP 실질 원계열 전년 동기비에 대해서는 TFI만 고려한 경우 역시 월별자료분석과 같이 슬라이스 수를 10으로 하여 SIR과 DR을 이용하여 차원 축소를 하였다. 다만 MA도 고려하는 경우 자료의 수가 16개로 매우 적기 때문에 분기별에서는 MA자료가 사용되지는 않는다.

차원축소 시 월별, 분기별 예측 목표 변수들과 TFI와 MA 간에 시차효과(lag effect)가 존재할 수 있기 때문에, 월별 자료의 경우는 시차로 최대 3개월의 시차를 고려한 반면, 분기별 자료의 경우는 최대 1분기 시차만 고려하였다. 이에 대한 이유는 텍스트 자료의 경우 공식 통계보다 선행은 하지만, 뉴스와 리포트 기반이기에 긴 시차효과는 없으리라 판단했기 때문이다. 최적 시차는 시차별 차원축소된 변수와 다섯 개의 예측 목표 변수 간의 선형회귀 모형을 적합하여 가장 높은 adjust R^2 값을 보이는 시차를 검토하여 분석에 활용하였다. 월별 분석에서는 선행지수 순환 변동치, BSI 전산업 매출실적과 BSI 전산업 매출전망에 대해 SIR을 이용하여 TFI만 차원축소를 한 경우 각각 2, 3와 0이 최적 lag이었고, DR을 적용하면 각각 2, 1와 3이었다. 반면 seeded method로 MA만 축소를 하는 경우에는 3, 1과 2이다. 마지막으로 seeded method로 TFI와 MA를 축소한 경우에는 선행지수 순환 변동치, BSI 전산업 매출실적과 BSI 전산업 매출전망에 대해 각각 3, 1와 1이다. 분기별

Table 1: Performance of monthly variable predictions using only TFI

Target	Model				
	HP λ	Data	Dimension reduction	Model	MAE
Cycle of Composite Leading Ind.	0.5	Offi. Stat. + TFI	Directional Regression	LR / Lasso	0.112
BSI All Industry Sales	2.0	Offi. Stat. + TFI	Not applied	LR	2.285
BSI All Industry Future Sales	0.5	Offi. Stat. + TFI	Directional Regression	Lasso	2.241

Table 2: Performance of monthly variable predictions using MA exclusively and in combination with TFI

Target	Model				
	HP λ	Data	Dimension reduction	Model	MAE
Cycle of Composite Leading Ind.	2.0	Offi. Stat. + MA	seeded method	Lasso	0.164
BSI All Industry Sales	Not applied	Offi. Stat. + TFI + MA	Not applied	Lasso	2.111
BSI All Industry Future Sales	Not applied	Offi. Stat.	Not applied	AR	1.772

Table 3: Performance of quarterly variable predictions using only TFI

Target	Model				
	HP λ	Data	Dimension reduction	Model	MAE
Real GDP SA QoQ	0.25/0.5	Offi. Stat. + TFI	SIR	Lasso	0.404
Real GDP YoY	0.25	Offi. Stat. + TFI	Not applied	Lasso	0.922

분석에서는 GDP 실질 SA전기비와 GDP 실질 원계열 전년 동기에 대해 SIR을 이용하여 TFI의 차원축소 시 최적 lag는 각각 0과 1인 반면, DR을 이용하면 0과 0이었다.

예측 모형으로서는 시계열분석에서 가장 고전적으로 널리 사용되는 자기회귀(autoregressive regression; AR)모형을 적합하고 BIC (Bayesian information criterion)를 기준으로 가장 적절한 p 를 분석하였다. 이를 바탕으로 AR모형을 기준으로 AR모형에 차원 축소된 텍스트 자료 그리고 AR모형에 차원 축소 전의 원 텍스트 자료를 고려한 세 가지 예측 모형을 적합하였다. AR모형에 텍스트 자료를 함께 고려한 예측 모형으로 차원 축소된 텍스트 자료를 이용한 경우에는 선형회귀 모형을 적합하였고, 원래의 텍스트 자료를 이용한 경우에는 Lasso 회귀를 적합하였다. 모형의 예측력을 비교하기 위하여, rolling window를 활용한 one-step ahead cross-validation을 통하여 테스트 자료의 에러를 확인하였다. 이에 대해 다음과 같이 정리할 수 있다. 월별 자료의 경우는 2020년 12월까지의 자료를 학습 데이터셋으로, 이후 2021년 1월부터 2022년 12월까지 2년치 자료를 검증 데이터셋으로 사용하였다. 반면에 자료의 수로 인해 분기별 자료의 경우는 2021년 4분기까지의 자료를 학습 데이터셋으로, 이후 2022년의 4분기자료를 검증 데이터셋으로 사용하였다.

네트워크 자료의 분석에 대해서는 이후의 절에서 상세히 설명하기로 한다.

5.2. 월별 텍스트 지표 분석

월별 자료의 예측 목표 변수는 선행지수 순환 변동치, BSI 전산업 매출실적, BSI 전산업 매출전망이다. 앞에서 언급된 분석 방법을 적용하여 MAE를 계산하였다.

먼저 TFI만을 이용한 분석을 살펴보고, 각 예측 목표 변수별 최적의 결과가 Table 1에 제시되어 있고, 상세한 결과는 Table 5에 정리되어 있다.

Table 1을 살펴보면, 우선 공식통계만 사용하기 보다는 TFI의 자료를 사용하는 것이 예측력을 높여주고 있다. 또한 HP filtering을 이용한 변동성 안정화가 예측을 보다 정확하게 해준다. 그리고 선행지수 순환 변동치와 BSI 전산업 매출전망에 대해서는 차원축소를 하는 것이 원자료를 사용하는 것보다 예측력을 높여주는 반면,

Table 4: Performance comparisons in MAE for monthly variables

Target	HP λ	AR	Model			
			Full MA		Dim. reduced MA	
			Linear Reg.	Lasso Reg.	Seed-Linear Reg.	Seed-Lasso Reg.
Cycle of Composite Leading Ind.	NA	0.167	1.186	0.242	0.185	0.186
	0.5		4.270	0.176	0.198	0.213
	1.0		1.457	0.180	0.191	0.186
	1.5		6.303	0.189	0.180	0.173
	2.0		1.454	0.185	0.174	0.164
BSI All Industry Sales	NA	2.642	73.228	2.959	5.172	5.166
	0.5		235.864	3.134	4.047	4.153
	1.0		40.759	2.375	3.73	3.766
	1.5		416.203	2.396	3.613	3.779
	2.0		77.485	2.444	3.613	3.619
BSI All Industry Future Sales	NA	1.772	141.230	2.915	3.244	3.404
	0.5		42.587	3.586	3.552	3.674
	1.0		45.885	2.215	2.348	2.458
	1.5		64.845	2.141	2.447	2.491
	2.0		60.346	2.111	2.575	2.502
Target	HP λ	AR	Model			
			TFI + MA		Dim. reduced TFI + Dim. reduced MA	
			Linear Reg.	Lasso Reg.	Seed-Linear Reg.	Seed-Lasso Reg.
Cycle of Composite Leading Ind.	NA	0.167	1.845	0.189	0.185	0.186
	0.5		4.551	0.146	0.198	0.214
	1.0		2.001	0.190	0.192	0.186
	1.5		2.990	0.199	0.181	0.180
	2.0		7.710	0.193	0.174	0.168
BSI All Industry Sales	NA	2.642	135.319	2.111	5.173	5.167
	0.5		329.866	2.270	4.041	4.145
	1.0		191.953	2.559	3.807	3.847
	1.5		23.694	2.784	3.641	3.813
	2.0		32.270	3.328	3.631	3.639
BSI All Industry Future Sales	NA	1.772	38.651	2.159	3.247	3.406
	0.5		63.975	2.336	3.547	3.669
	1.0		33.757	2.536	2.498	2.563
	1.5		237.391	2.705	2.459	2.501
	2.0		84.724	2.502	2.590	2.651

BSI 전산업 매출실적에서는 차원축소를 하지 않은 원자료를 사용하는 것이 가장 높은 예측력을 보여주고 있다.

MA가 사용되는 월별 자료 분석에 대한 최적 결과는 Table 2에 그리고 보다 상세한 결과는 Table 4에 정리되어 있다.

Table 2를 살펴보면, HP filtering을 이용한 변동성 안정화는 선행지수 순환 변동치에서만 최적의 예측력을 제시하고 있다. 그리고 텍스트 자료의 예측력 강화 측면에서는 세 예측 목표 변수에 대해 결과가 제시되고 있다. 선행지수 순환 변동치에서는 MA만 사용하는 것이, BSI 전산업 매출실적에서는 TFI와 MA를 모두 사용하는 것이, 마지막으로 BSI 전산업 매출전망에서는 모두 사용하지 않는 것이 최적의 예측력을 나타낸다.

보다 상세한 모형 평가 결과는 Table 5과 4에 기술되어 있다. 선행지수 순환변동치의 경우 원 MA 또는 차원축소된 MA를 사용하는 경우, 그리고 linear regression 또는 Lasso regression을 사용하는 경우의 예측성능

Table 5: Performance comparisons in MAE for monthly variables

Target	HP λ	AR	Model									
			TFI		One dimensional TFI				Two dimensional TFI			
			LR	Lasso	SIR LR	SIR Lasso	DR LR	DR Lasso	SIR LR	SIR Lasso	DR LR	DR Lasso
Cycle of Composite Leading Ind.	NA	0.124	0.124	0.124	0.126	0.127	0.119	0.119	0.122	0.121	0.124	0.124
	0.5		0.122	0.121	0.123	0.123	0.112	0.112	0.128	0.127	0.116	0.115
	1.0		0.139	0.134	0.130	0.129	0.114	0.113	0.124	0.123	0.117	0.115
	1.5		0.155	0.145	0.134	0.133	0.114	0.113	0.117	0.118	0.141	0.142
	2.0		0.151	0.142	0.129	0.130	0.127	0.127	0.116	0.116	0.141	0.142
BSI All Industry Sales	NA	2.618	2.689	2.462	2.517	2.523	3.029	3.040	2.586	2.591	2.696	2.701
	0.5		2.678	2.44	2.708	2.717	3.008	3.017	2.898	2.908	2.741	2.752
	1.0		2.567	2.446	2.826	2.825	2.782	2.784	3.085	3.083	2.841	2.838
	1.5		2.571	2.463	2.646	2.650	2.746	2.749	2.906	2.888	3.151	3.133
	2.0		2.285	2.328	2.685	2.665	2.723	2.723	2.765	2.742	3.640	3.590
BSI All Industry Future Sales	NA	2.554	2.845	2.564	2.561	2.571	2.546	2.574	2.443	2.474	2.468	2.453
	0.5		2.900	2.436	2.565	2.566	2.568	2.581	2.540	2.561	2.271	2.241
	1.0		2.765	2.487	2.724	2.729	2.623	2.639	3.144	3.125	2.949	2.897
	1.5		2.582	2.446	2.542	2.571	2.611	2.618	3.120	3.088	3.869	3.800
	2.0		2.381	2.301	2.566	2.589	2.598	2.595	2.84	2.732	4.187	4.097

차이는 크지 않은 반면, 변동성 안정화 λ 값의 변화에 따른 예측성능 차이는 상대적으로 크게 나타난다. 또한 원 TFI와 원 MA를 linear regression으로 분석하는 경우 차원의 저주로 인해 예측오차가 크게 높아지는 것을 알 수 있다. 한편 BSI 전산업 매출실적 및 전망의 경우에는 모형 및 변동성 안정화 λ 값 변화에 따른 예측 성능차이가 대체로 높게 나타나는 것으로 분석되었다.

월별 자료분석에 있어 이렇게 차이가 나타나는 가장 큰 이유로 자료의 크기를 들 수 있다. TFI만 사용하는 경우에는 2005년부터의 월별 자료를 사용할 수 있는 반면, MA를 고려하면 2019년도의 자료만 사용하게 된다. 이러한 자료의 크기는 결과의 차이에 직접적인 영향을 줄 수 있다. 특히 변동성 안정화의 경우는 긴 시간에 걸친 자료가 있을 때 그 효과가 나타남을 확인할 수 있고, 자료의 차원 축소 또한 많은 자료에서 보다 효과적인 차원 축소가 가능함을 알 수 있다.

5.3. 분기별 텍스트 지표 분석

분기별 자료의 예측 목표 변수는 GDP 실질 SA전기비와 GDP 실질 원계열 전년 동기비이며, 두 경제 지표를 예측하기 위해 월별 TFI 텍스트 지표만을 고려한다. MA 자료가 고려되지 않은 이유는 앞에서 언급하였듯이 2019년부터의 분기별 자료의 수가 매우 적어 분석의 결과에 대해 신뢰할 수 없기 때문이다. 예측에 대한 최적 결과는 Table 3에 정리되어 있고, 상세 결과는 Table A.4을 참고하길 바란다.

Table 3에서 확인 할 수 있듯이, 두 예측 목표 변수에 대해 변동성 안정화는 더 나은 예측력을 확보하게 해준다. 또한 TFI 자료의 활용은 예측력을 높여줄 수 있다. 그리고 GDP 실질 SA 전기비에 대해서는 SIR을 이용한 차원 축소가 예측력 향상에 효과적인 반면 GDP 실질 원계열 전년 동기비의 경우에는 원자료의 TFI를 사용하는 것이 최적의 결과를 제공함을 확인할 수 있다.

5.4. 네트워크 자료 분석

애널리스트 리포트 기반으로 추정된 산업별 유사도 네트워크 자료(cofactor)는 MA자료 중 변수명 1번에서 39번까지의 변수들에 대한 관계를 분기별 그리고 연별 계산한 39×39 의 네트워크 행렬자료이다. 이 자료는 예측 목표 변수의 예측에 직접적으로 사용할 수 없다. 우선 편의상 MA의 1-39변수의 자료를 T1이라고 한다.

우선 각각의 분기별 네트워크 행렬을 spectral decomposition을 이용하여 분해 한 후 가장 큰 고유치와 상

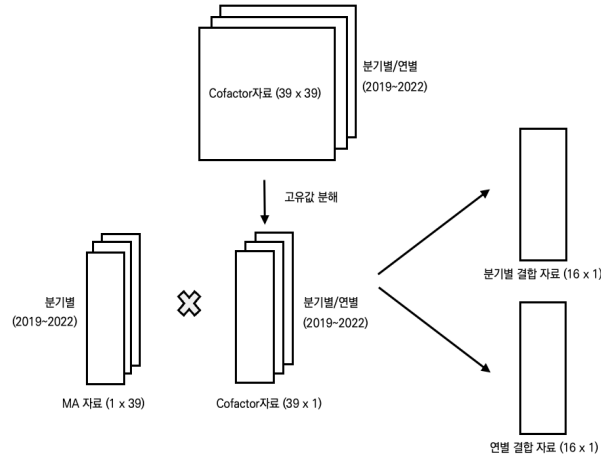


Figure 2: Schematic illustration of integrating network data from analyst reports with MA text data.

Table 6: Correlation between network data and dimension reduced MA data

Dimension reduction	Network data	Corr. Coef.
Real GDP SA QoQ dim. red. w/ seeded method	Quarterly (QN)	-0.625
	Annual (AN)	-0.648
Real GDP YoY dim. red. w/ seeded method	Quarterly (QN)	-0.619
	Annual (AN)	-0.619
1st PCA of MA	Quarterly (QN)	0.156
	Annual (AN)	-0.443

응하는 고유벡터를 구한다. 이 고유 벡터를 l_i 라고 정의한다. 이후 각 분기에 해당되는 T1의 관측값 $T1_i$ 와 위에서 구한 고유 벡터를 l_i 를 곱하여, 즉 $l_i^T T1_i$, 새로운 관측값을 생성한다. 이렇게 계산한 관측값을 QN_i 라고 하자.

이와 동일하게 연별 네트워크 행렬 역시 spectral decomposition을 이용하여 고유벡터를 구한 후 해당 연도에 대응되는 T1의 분기별 관측값을 곱하여 또 다른 관측값을 생성한다. 이 관측값을 AN_i 라고 하자. 분기별 네트워크 MA 관측값인 QN_i 와 연별 네트워크 MA 관측값은 AN_i 는 분석 가능한 형태의 자료이다. 이 두 자료에 대한 계산 과정은 Figure 2를 참고하길 바란다.

QN_i 와 AN_i 가 실제로 MA의 자료와 어떠한 관계성을 가지고 있나 파악하기 위해 분기별 MA 자료를 GDP 실질 SA전기비와 GDP 실질 원계열 전년 동기비에 대해 seeded method을 적용하여 차원 축소된 자료와 분기별 MA 자료에 대한 첫번째 주성분에 대한 상관분석을 실시하였다. 상관 분석을 위하여 피어슨의 상관계수를 계산하였고, 이는 Table 6에 정리되어 있다.

Table 6를 살펴보면 QN과 AN은 예측 목표 변수를 고려하지 않은 MA의 첫 번째 주성분보다는 예측 목표 변수를 고려하여 축소된 변수들과 월등히 높은 상관관계를 갖고 있음을 확인할 수 있다. 또한 QN과 AN은 모두 seeded method로 축소된 MA 자료와 높은 상관관계를 갖고 있기 때문에, QN과 AN은 바로 해당 변수의 분석에 사용될 수 있음을 파악할 수 있어, 네트워크 자료의 실제적 유용성에 대한 가능성을 확인할 수 있다.

6. 결론 및 제언

본 연구는 노이즈 필터링 기법과 차원축소 등의 방법을 이용하여 텍스트 지표의 정상화에 대해 검토하고 실증분석을 통해 텍스트 지표의 활용가능성을 높이고자 하였다. 노이즈 필터링의 방법으로 Hodrick과 Prescott (1997)이 제안한 필터를 이용하여 뉴스 및 애널리스트 보고서 기반의 텍스트 지표에 다양한 조절 모수 값들을 적용함으로써 노이즈를 제거하였다. 이후 필터링 된 지표를 이용하여 월별 선행지수 순환 변동치, BSI 전산업 매출실적, BSI 전산업 매출전망 그리고 분기별 GDP 실질 SA전기비와 GDP 실질 원계열 전년 동기비 예측 목표 변수에 대해 회귀분석에서 비모수적으로 차원을 축소하는 충분차원축소 방법론 중에서 sliced inverse regression (Li, 1991), directional regression (Li와 Wang, 2007), seeded method (Cook 등, 2007)을 적용 하였다.

분석 결과 월별과 분기별 변수 모두 자료의 수가 많은 경우에 있어서는 텍스트 지표의 사용과 이에 대한 노이즈 필터링이 예측의 정확도를 높이는 것으로 나타나고 있다. 그리고 텍스트 지표 자료의 차원 축소를 함에 따라 예측을 보다 더 정확하게 할 수 있음을 제시하고 있다. 하지만 자료의 수가 상대적으로 적은 경우 노이즈 필터링과 차원 축소의 효과가 기대와는 달리 미비하였다. 이것이 시사하는 바는 자료가 충분할 경우 텍스트 지표의 사용은 노이즈 필터링과 차원 축소를 통해 예측의 정도를 높일 수 있다는 것이다.

차원 축소의 또 다른 장점으로 두 텍스트 지표 자료를 저차원으로 변환함으로써 예측 목표 변수들과 차원 축소된 지표들에 대한 시차 효과를 파악하는 것이 원지표를 고려하는 것보다 용이하다.

또한 애널리스트 지표의 네트워크 행렬 자료는 그 자체로 분석에 직접적인 사용은 어려우나 행렬의 차원 축소결과와 애널리스트 텍스트 지표와의 결합을 통해 실제 분석에서 사용되는 차원 축소된 애널리스트 텍스트 지표와 높은 상관 관계를 가지고 있어 그 자체만으로 의미 있는 자료가 될 수 있음을 분석을 통해 확인하였다.

본 연구에 대한 사후 연구 문제로 다음을 제안하고자 한다. 뉴스 기반 지표와 애널리스트 보고서 기반 지표에 대한 차원 축소를 개별적으로 진행하였다. 하지만 애널리스트 보고서 기반 지표가 시간에 걸쳐 충분히 누적 된다면 두 지표를 동시에 축소하는 방법이 고려될 수 있고, 그렇다면 두 지표 간의 상관성이 차원 축소하는 과정에 고려될 수 있어 더 바람직한 차원 축소가 이루어질 수 있으리라 기대할 수 있다. 그리고 두 텍스트 자료의 차원 축소는 노이즈 필터링 이후의 자료에 대한 차원 축소이기 때문에 sliced inverse regression이나 directional regression 등과 같은 선형 공간의 차원 축소 방법보다는 비선형 공간에서의 차원 축소 방법이 더 적합할 수 있다.

끝으로 많은 결측치를 선형 보간법으로 대체를 하였는데, 다른 보간법을 사용한 후 예측 모형을 개발하여 정확도를 비교함으로써, 해당 텍스트 자료에 보다 적합한 보간법을 찾아보는 것도 의미 있는 연구가 될 것이다.

Appendix A:

Table A.1: TFI text indices

No.	Sector	No.	Sector	No.	Sector
1	Construction	6	Production	11	Automobiles
2	Job Search	7	Shipbuilding	12	Gov. Expenditure
3	Retail & Wholesale	8	Facility Inv.	13	Stock Price
4	Price Expectation	9	World Trade	14	House Price
5	Semiconductor	10	Unemployment	15	Recruitment

Table A.2: MA text indices

No.	Sector	No.	Sector	No.	Sector
1	Primary Metals	15	Real Estate	29	Printing/Recorded Media
2	Furniture	16	Non-Metallic Minerals	30	Automobiles
3	Leather/Goods/Shoes	17	Business Support/Lease	31	Electricity/Gas/Steam
4	Construction	18	Petroleum/Coke	32	Electrical Equipment
5	Rubber/Plastics	19	Textiles	33	Scientific/Technical
6	Education Services	20	Accommodations	34	Communication Equipment, etc.
7	Metal Processing	21	Food	35	Information Services
8	Finance	22	Fisheries	36	Shipbuilding
9	Other Personal Serv.	23	Arts/Sports/Leisure	37	Pulp/Paper
10	Other Machine./Equip.	24	Transport/Warehousing	38	Sewage/Waste Management
11	Other Manufact.	25	Beverages	39	Chemicals/Products
12	Agriculture	26	Medical/Precision Equip.	40	Computers
13	Wholesale/Retail	27	Medical/Pharma.	41	Other Transportation
14	Wood/Forestry	28	Clothing/Fur		

Table A.3: Network indices

No.	Sector	No.	Sector	No.	Sector
1	Primary Metals	14	Wood/Forestry	27	Medical/Pharma.
2	Furniture	15	Real Estate	28	Clothing/Fur
3	Leather/Goods/Shoes	16	Non-Metallic Minerals	29	Printing/Recorded Media
4	Construction	17	Business Support/Lease	30	Automobiles
5	Rubber/Plastics	18	Petroleum/Coke	31	Electricity/Gas/Steam
6	Education Services	19	Textiles	32	Electrical Equipment
7	Metal Processing	20	Accommodations	33	Scientific/Technical
8	Finance	21	Food	34	Communication Equipment
9	Other Personal Serv.	22	Fisheries	35	Information Services
10	Other Machine./Equip.	23	Arts/Sports/Leisure	36	Shipbuilding
11	Other Manufact.	24	Transport/Warehousing	37	Pulp/Paper
12	Agriculture	25	Beverages	38	Sewage/Waste Management
13	Wholesale/Retail	26	Medical/Precision Equip.	39	Chemicals/Products

Table A.4: Performance comparisons in MAE for quarterly variables

Target	Model											
	HP λ	AR	TFI		One dimensional TFI				Two dimensional TFI			
			LR	Lasso	SIR LR	SIR Lasso	DR LR	DR Lasso	SIR LR	SIR Lasso	DR LR	DR Lasso
Real GDP SA QoQ	NA	0.523	1.348	1.344	0.682	0.453	0.650	0.614	0.832	0.434	0.631	0.579
	0.5		1.373	1.346	0.646	0.464	0.637	0.597	0.796	0.404	0.604	0.572
	1.0		1.442	1.396	0.536	0.495	0.627	0.578	0.734	0.404	0.588	0.543
	1.5		1.455	1.433	0.502	0.508	0.624	0.577	0.723	0.626	0.570	0.539
	2.0		1.464	1.440	0.479	0.519	0.607	0.565	0.719	0.630	0.554	0.559
Real GDP NSA YoY	NA	1.252	1.259	0.938	1.375	1.292	1.106	1.066	1.431	1.323	1.043	1.047
	0.5		1.225	0.922	1.349	1.273	1.104	1.072	1.412	1.297	0.965	0.968
	1.0		1.062	0.925	1.312	1.270	1.393	1.323	1.338	1.308	1.763	1.696
	1.5		1.004	0.949	1.319	1.277	1.390	1.321	1.319	1.277	1.39	1.321
	2.0		0.973	0.980	1.332	1.286	1.383	1.315	1.407	1.341	1.799	1.725

References

- Chen CC, Huang HH, Huang YL, and Chen HH (2021, October). Constructing noise free economic policy uncertainty index. In *Proceedings of the 30th ACM International Conference on Information & Knowledge Management*, Gold Coast, Queensland, Australia, 2915–2919.
- Cook RD, Li B, and Chiaromonte F (2007). Dimension reduction in regression without matrix inversion, *Biometrika*, **94**, 569–584.
- Hall P and Li KC (1993). On almost linearity of low dimensional projections from high dimensional data, *Annals of Statistics*, **21**, 867–889.
- Hamilton JD (2018). Why you should never use the Hodrick-Prescott filter, *Review of Economics and Statistics*, **100**, 831–843.
- Hodrick R and Prescott EC (1997). Postwar U.S. business cycles: An empirical investigation, *Journal of Money, Credit and Banking*, **29**, 1–16.
- Kim Y, Ahn YH, Yoo JK, and Kwon O (2017). Verifying identities of plant-based multivitamins using phytochemical fingerprinting in combination with multiple bioassays, *Plant Foods for Human Nutrition*, **72**, 288–293.
- Lee K, Choi Y, Um HY, and Yoo JK (2019). On fused dimension reduction in multivariate regression, *Chemometrics and Intelligent Laboratory Systems*, **193**, 103828.
- Li B (2018). *Sufficient Dimension Reduction Methods and Applications with R*, Chapman and Hall/CRC, London, UK.
- Li B and Wang S (2007). On directional regression for dimension reduction, *Journal of the American Statistical Association*, **102**, 997–1008.
- Li B, Zha H, and Chiaromonte F (2005). Contour regression: A general approach to dimension reduction, *Annals of Statistics*, **33**, 1580–1616.
- Li KC (1991). Sliced inverse regression for dimension reduction, *American Statistical Association*, **86**, 316–327.
- Ravn MO and Uhlig H (2002). On adjusting the Hodrick-Prescott filter for the frequency of observations, *Review of Economics and Statistics*, **84**, 371–376.
- Li, B., Zha, H., and Chiaromonte, F. (2005). Contour regression: A general approach to dimension reduction, *Annals of Statistics*, **33**, 1580–1616.
- Seo B (2023a). Econometric Forecasting Using Ubiquitous News Text: Text-enhanced Factor Model, *Bank of Korea WP*, **10**, 1–48.
- Seo B (2023b). Industry analysis using AI algorithms: Text analysis of securities research reports, *Bank of Korea Issue Note*, **5**, 1–26.
- Seo B, Lee Y, and Cho H (2024). Measuring News Sentiment of Korea Using Transformer. *Korean Economic Review*, **40**, 149–176.
- Sung T and Si G (2020). *Research Methodology*, (3rd ed), Hakjisa, Seoul.
- Um HY, Won S, An H, and Yoo JK (2018). Case study: Application of fused sliced average variance estimation to near-infrared spectroscopy of biscuit dough data, *The Korean Journal of Applied Statistics*, **31**, 835–842.
- Yoo JK (2019). Analysis of microarray right-censored data through fused sliced inverse regression, *Scientific Reports*, **9**, 15094.

Received July 7, 2023; Revised October 5, 2023; Accepted October 7, 2023

노이즈 필터링과 충분차원축소를 이용한 비정형 경제 데이터 활용에 대한 연구

유재근^a, 박유진^a, 서범석^{1,b}

^a이화여자대학교 통계학과; ^b한국은행 경제모형실

요 약

본 연구는 노이즈 필터링과 차원축소 등의 방법을 이용하여 텍스트 지표의 정상화에 대해 검토하고 실증 분석을 통해 동 지표의 활용가능성을 제고할 수 있는 후처리 과정을 탐색하고자 하였다. 실증분석에 대한 예측 목표 변수로 월별 선행지수 순환 변동치, BSI 전산업 매출실적, BSI 전산업 매출전망 그리고 분기별 실질 GDP SA전기비와 실질 GDP 원계열 전년동기비를 상정하고 계량경제학에서 널리 활용되는 Hodrick and Prescott 필터와 비모수 차원축소 방법론인 충분차원축소를 비정형 텍스트 데이터와 결합하여 분석하였다. 분석 결과 월별과 분기별 변수 모두에서 자료의 수가 많은 경우 텍스트 지표의 노이즈 필터링이 예측 정확도를 높이고, 차원 축소를 적용함에 따라 보다 높은 예측력을 확보할 수 있음을 확인하였다. 분석 결과가 시사하는 바는 텍스트 지표의 활용도 제고를 위해서는 노이즈 필터링과 차원 축소 등의 후처리 과정이 중요하며 이를 통해 경기 예측의 정도를 높일 수 있다는 것이다.

주요용어: 빅데이터, 텍스트 자료, 차원축소, 예측모형

본 연구는 한국은행의 연구용역지원을 받아 수행되었습니다.

유재근과 박유진은 2023년도 정부(과학기술정보통신부)의 재원으로 한국연구재단의 지원을 받아수행된 기초연구사업 임 (RS-2023-00240564 and RS-2023-00217022).

¹교신저자: (04531) 서울특별시 중구 남대문로 39 (남대문로3가), 한국은행 경제모형실. E-mail: bsseo@bok.or.kr