

글의 수준을 평가하기 위한 글의 응집성 측정

조용구 광주교육대학교 강사

- I. 서론
- II. 이론적 배경
- III. 연구 방법
- IV. 결과 및 논의
- V. 결론

I. 서론

『곰 세 마리』라는 영국의 전래 동화가 있다. 이 동화의 줄거리는 다음과 같다. ‘골디락스’라는 이름의 금발머리 소녀가 있었다. 이 소녀는 숲 속을 헤매다가 우연히 곰 가족의 오두막에 들어갔다. 당시 곰 가족은 뜨거운 죽이식을 동안 산책을 나간 상태였다. 소녀는 아빠 곰의 죽은 너무 뜨거워서 먹지 않고, 엄마 곰의 죽은 너무 차가워서 먹지 않았다. 대신 온도가 적당한 아기 곰의 죽을 맛있게 먹었다.

이 동화에서 유래하여, 금발머리라는 뜻의 골디락스(Goldilocks)는 다양한 분야에서 사용되었다. 이 용어는 주로 경제 분야에서 사용되었는데, 경제 성장이 계속되지만 물가가 상승하지 않는 상태를 이른다. 즉, 너무 뜨겁지도 않고 차갑지도 않은 호황기의 경제 상황을 일컫는다. 이 용어는 읽기 교육에서도 사용된다. 독자에게 너무 어렵지도 않고 너무 쉽지도 않은 적정 수준의 글을 제공하는 것을 ‘골디락스 원리’라고 한다.

적정 수준의 글을 제공하기 위해 읽기 교육 연구자들은 많은 노력을 해왔다. 이러한 노력은 어떠한 언어적 변수가 글을 어렵게 하는지 찾는 것으로

부터 시작한다. 글을 어렵게 하는 언어적 변수에는 수백 가지가 있는데, ‘단어의 길이’, ‘단어의 친숙성’, ‘문장의 길이’, ‘문장의 구조’, ‘글의 응집성’ 등이 그 예이다. 연구자들은 글에서 이러한 언어적 변수를 측정 한 후에, ‘글의 수준(text difficulty)’을 평가하고자 하였다.

그런데 다른 언어적 변수에 비해 ‘글의 응집성(text coherence)’은 상대적으로 주목을 받지 못하였다. 응집성이 주목받지 못한 이유는 객관적으로 측정하기 어렵기 때문이다. ‘단어의 길이’ 및 ‘문장의 길이’와 같은 변수는 표층 수준의 변수이고, 따라서 측정하기가 비교적 쉽다. 하지만 응집성은 심층 수준의 변수이고, 따라서 측정하기가 어렵다. 이러한 이유로 1980년대에 ‘글의 응집성’ 및 ‘글의 구조’와 같은 심층 수준의 변수를 측정하려는 시도(예, Meyer, 1982)가 있었지만, 소수에 그쳤으며 글의 수준을 평가하는 데 잘 활용되지 못하였다.

하지만 글의 응집성은 글의 수준에 영향을 주는 중요한 변수이다. 응집성은 글의 여러 자질 중 핵심적인 자질이기 때문이다(김명순, 2010: 182). 일반적으로 글의 응집성이 높을수록, 그 글은 읽기 쉽다. 반대로 응집성이 낮을수록, 글은 읽기 어렵다. 다른 변수가 동일하다면, 응집성이 낮은 글이 어려운 글이 된다. 따라서 저학년에서 고학년으로 갈수록, 글의 응집성이 낮아지는 것이 이상적이다. 응집성을 효율적으로 측정할 수 있다면, 글의 수준 평가에서 고려하지 못했던 중요한 측면을 파악할 수 있게 된다.

이에 따라, 본 연구의 목적은 글의 수준을 평가하기 위해 글의 응집성을 측정하는 데 있다. 글의 응집성을 측정하는 기법으로서 ‘잠재의미분석’을 사용한다. 아래에서는 응집성과 잠재의미분석에 대하여 간략히 기술한 후에, 초등학교 국어 교과서에 실린 글을 시험적으로 분석한다.

II. 이론적 배경

1. 글의 응집성

응집성(coherence)은 글에 포함되어 있는 내용들 간의 의미적인 연결 관계를 말한다(이재승, 2003: 94). 하나의 글은 이러한 연결 관계를 통하여 전반적인 의미를 전달한다. 독자가 글을 잘 이해하기 위해서는 머릿속에 글의 내용에 대한 연결된 표상을 형성해야 한다. 따라서 독해는 응집성의 영향을 크게 받는다. 응집성은 흔히 결속 구조(cohesion)와 함께 논의된다. 결속 구조는 글의 표층에서 연결 관계를 나타내는데, 접속어 등이 이에 해당한다. 이에 반해 응집성은 글의 심층에서 연결 관계를 나타낸다. 이들이 엄격하게 구분되기 보다는 상호작용하는 것으로 볼 수 있다.

응집성은 글의 수준에 영향을 준다. 일반적으로 응집성이 높을수록, 글이 쉬워진다. Kintsch et al.(1975: 201-205)은 글의 응집성이 높을수록, 읽기 시간이 감소하고 회상이 증가한다는 것을 발견하였다. 즉, 독자들은 응집성이 높은 글을 더 쉽게 이해할 수 있었다. Todorascu et al.(2013: 16-17)은 응집성에 관한 몇몇 하위 요소가 글의 수준과 유의미하게 상관된다는 것을 발견하였다. McNamara et al.(2010: 4-11)은 응집성에 관련된 12개의 선행 연구를 검토하였다. 그 결과, 높은 응집성은 대부분의 경우에 독해를 향상시키는 경향이 있었다. 특정 화제에 대해 배경지식이 낮은 독자들은 응집성이 높은 글을 더 잘 이해하였다. 또한 능숙한 독자와 미숙한 독자 모두 응집성이 높은 글을 더 잘 이해하였다.

그동안 응집성은 주로 교사나 전문가의 질적 판단에 의해 평가되었다. 예를 들어, 교사들은 1~5점의 리커트 척도로 응집성의 정도를 평가하였다. 하지만 질적 평가의 단점은 평가 결과가 주관적이라는 것이다. 동일한 글에 대하여 평가자들은 서로 다른 판단을 내릴 수 있다. 즉, 어떤 사람은 응집

성이 낮은 글로 판단하고, 다른 사람은 응집성이 높은 글로 판단할 수 있다.

응집성에 대한 질적 평가를 대체할 수 있는 방법은 양적 평가이다. 전통적인 양적 평가 방법은 ‘명제 분석’이었다. 명제(proposition)는 술어와 논항으로 이루어진다. 예를 들어, ‘영희가 철수에게 책을 주었다.’라는 문장을 명제로 분석하면 다음과 같다.

주다(행위주: 영희, 대상: 책, 목표: 철수)

이 명제에서 ‘주다’는 술어이고, ‘영희’, ‘책’, ‘철수’는 논항이다. 이 문장 다음에 ‘영희가 도서관에서 나왔다.’라는 문장이 이어진다면, ‘영희’라는 논항이 반복된다. 이러한 논항 반복은 글의 응집성에 기여한다. 이처럼 글을 명제 단위로 분석한 후에, ‘논항 반복의 비율’을 계산하여 응집성을 측정할 수 있다. 논항 반복의 비율이 클수록, 응집성이 높다. 그러나 이 방법의 단점은 글을 명제로 분석하는 데 시간이 오래 걸리고, 상당한 전문성이 요구된다는 점이다(Foltz et al., 1998: 3). 이에 따라 분석할 수 있는 글의 크기에 제한이 생기고, 사례 연구의 성격을 지니게 된다. 명제 분석은 응집성을 측정할 수 있는 반면에, 활용상에 단점을 지니고 있다.

2. 잠재의미분석

명제 분석을 보완할 수 있는 또 다른 양적 평가 방법은 ‘잠재의미분석’이다. 잠재의미분석(Latent Semantic Analysis, 이하 LSA)은 의미 관계를 추출하고 표상하기 위한 자동화된 수학적 기법이다(Landauer et al., 1998: 8). LSA는 1980년대 후반에 정보 검색 기법으로 만들어졌으며, 그 이후에 인공지능, 인지과학, 교육학, 심리학 등의 분야로 확장되었다(Evangelopoulos et al., 2012: 7).

LSA의 핵심은 ‘상호 제약’을 활용하는 능력이다(Landauer, 2007: 13).

상호 제약의 간단한 예로 다음과 같은 방정식을 들 수 있다.

$$A + 2B = 8$$

$$A + B = 5$$

위의 방정식은 하나의 식만으로 A나 B의 값을 말할 수 없다. 두 식이 함께 쓰일 때, A는 2이고 B는 3이라고 말할 수 있다. LSA의 상호 제약도 이와 유사하다. LSA는 글과 단어의 의미를 다음과 같은 식으로 나타낸다.

$$\text{글 1 의미} = \Sigma(\text{단어 1 의미}, \text{단어 2 의미}, \dots \text{단어 n 의미})$$

위의 식에서 글의 의미는 단어의 의미의 합이다. 그리고 LSA는 말뭉치를 수집하여 분석하는데, 말뭉치에는 다수의 글이 포함되어 있다. 말뭉치 안에서 단어와 글은 서로의 의미를 제약한다. LSA는 단어와 글 사이에 이러한 상호 제약을 통하여 의미를 추출하고 표상한다.

LSA의 기본 가정은 여러 개의 문서를 가로질러 단어 사용의 패턴에서 “잠재적인” 구조가 있고, 통계적 기법으로 이러한 잠재적인 구조를 추정할 수 있다는 것이다(Foltz, 1996: 198). LSA가 ‘잠재 의미(latent semantic)’를 분석한다는 것은 이와 같은 심층의 구조를 분석한다는 뜻이다. 여기에서 문서는 단어가 발생하는 맥락을 나타내는 것으로서, 문장일 수도 있고 문단일 수도 있고 텍스트일 수도 있다. LSA는 단어들과 문서들 사이의 연관성을 분석함으로써, 유사한 맥락에서 사용되는 단어들이 의미적으로 연관된다는 것을 나타낸다.

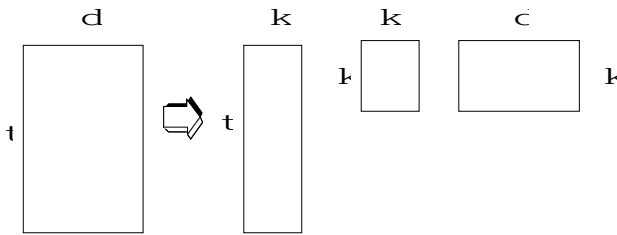
LSA는 크게 3단계의 과정을 거쳐 수행된다. 첫 번째 단계는 수집한 말뭉치를 행렬(matrix)로 만드는 것이다. 행렬의 일부를 예로 들면, 다음 <그림 1>과 같다.

Terms	Docs				
	1	2	3	4	5
가장자리/nc	0	1	0	0	0
가치/nc	0	0	0	1	0
각종/nc	0	0	0	1	0
갈대/nc	0	1	3	0	0
갈대숲/nc	0	0	1	0	0
강/nc	0	2	1	0	0
갯기미취/nc	0	1	0	0	0
갯메꽃/nc	0	1	0	0	0
갯벌/nc	1	2	0	6	0
갯질경이/nc	0	1	0	0	0

〈그림 1〉 $t \times d$ 행렬(일부)

위의 그림에서 행은 단어(terms)이고, 열은 문서(documents)이다. 여기에서 문서는 문단으로 설정하였다. 행렬에서 각각의 칸은 단어가 문단에서 발생하는 빈도를 나타낸다. 예를 들어, ‘가장자리’라는 단어는 문단 2에서 1번 발생한다. ‘갈대’는 문단 2에서 1번 발생하고, 문단 3에서 3번 발생한다. 행렬을 만든 후에, 빈도에 가중치를 부여한다. 가중치를 부여하면, 분석의 정확성을 향상시킬 수 있다(Quesada, 2007: 80). 가중치의 종류에는 로그 빈도, 역문서 빈도, 엔트로피 등 여러 가지가 있다.

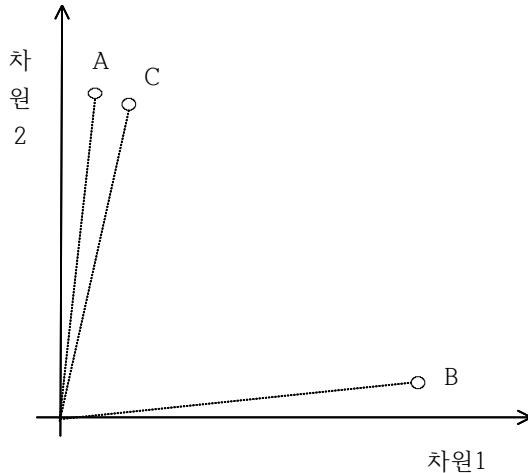
두 번째 단계는 행렬을 분해하는 것이다. 요인분석과 유사한 단일값분해(Singular Value Decomposition, 이하 SVD)를 사용한다. SVD를 간단하게 나타내면, 다음 〈그림 2〉와 같다(강범모, 2014: 6).



〈그림 2〉 $t \times d$ 행렬에 대한 SVD

위의 그림은 원래의 ‘ $t \times d$ ’ 행렬을 k 차원으로 분해한 것을 나타낸다. 이러한 k 차원의 의미 공간은 LSA에서 사용하는 의미 구조의 기초가 된다 (Martin & Berry, 2007: 42). 원래의 행렬을 k 차원으로 축소함으로써, ‘잡음 (noise)’으로 불리는 쓸모없는 정보들이 제거된다. SVD를 통해 단어와 문서의 잠재적인 의미 구조를 포착할 수 있다. 각각의 단어와 문서는 k 차원의 의미 공간에서 벡터(위치)로 표상된다. 의미가 유사한 단어들은 한 문서 내에서 공기(co-occur)하지 않더라도, k 차원의 의미 공간에서 서로 가깝게 위치한다. 또한 의미가 유사한 문서들은 공통된 단어를 가지지 않더라도, 의미 공간에서 가깝게 위치한다.

세 번째 단계는 벡터를 조작해서 의미적 거리를 계산하는 것이다. 의미적 거리를 계산하기 위해서 ‘코사인 유사도(cosine similarity)’가 주로 사용된다. 예를 들어, 세 단어(A, B, C)의 벡터가 다음 <그림 3>과 같다고 하자.



<그림 3> 코사인 유사도

위 그림은 2차원의 공간을 예로 든 것이다. 단어 A와 단어 B가 이루는 각은 90° 에 가깝다. $\cos 90^\circ = 0$ 이므로, 단어 A와 단어 B의 코사인 유사도는 0에 가깝다. 단어 A와 단어 B는 2차원의 공간에서 서로 다른 의미를 가진다.

반면, 단어 A와 단어 C가 이루는 각은 0° 에 가깝다. $\cos 0^\circ = 1$ 이므로, 단어 A와 단어 C의 코사인 유사도는 1에 가깝다. 단어 A와 단어 C는 2차원의 공간에서 유사한 의미를 가진다. 코사인 유사도가 0에 가까울수록, 두 단어는 의미적으로 관련이 없는 단어들이다. 반대로 코사인 유사도가 1에 가까울수록, 두 단어는 의미적으로 관련이 깊은 단어들이다. 여기에서는 2차원의 공간을 예로 들었지만, 수백 차원의 공간에서도 동일한 설명이 가능하다. 또한 앞서 언급한 대로, 글의 의미는 단어 의미의 합으로 나타낼 수 있으므로, 단어들 사이의 코사인 유사도는 문장들, 문단들, 텍스트들 사이에도 그대로 적용될 수 있다.

이와 같은 방식으로 LSA는 언어의 많은 부분에서 인간처럼 할 수 있는 연산 모형이다(Landauer, 2007: 5). 여러 연구에서 LSA의 효율성이 입증되었다. 이 중 교육과 직접 관련된 것 몇 가지만 살펴본다. LSA는 선다형 어휘 검사에서 고등학생과 비슷한 점수를 받았다. 이 검사 문항은 ‘보기’로 제시된 단어와 의미상 가장 가까운 단어를 고르는 것이다. 또한 LSA는 학생들이 작성한 서술형 답안을 전문가와 유사하게 평가할 수 있었다. 모범 답안과의 코사인 유사도에 기초하여, 학생들의 답안에 적절한 점수를 부여할 수 있었다. 그리고 LSA는 글의 응집성을 효율적으로 측정할 수 있었다. LSA는 응집성을 측정할 수 있을 뿐만 아니라, 응집성이 현저하게 떨어지는 부분도 포착할 수 있다. 본 연구는 LSA의 다양한 활용가능성 중 글의 응집성을 측정하는 능력에 초점을 두었다.

III. 연구 방법

1. 분석 대상

본 연구에서는 초등학교 국어 교과서에 실린 글의 응집성을 측정하였

다. 2015개정 교육과정에 따른 3~4학년군 국어 교과서에 수록된 정보전달 글을 분석하였다. 3학년 1학기 (가), (나)와 4학년 1학기 (가), (나)를 분석하였는데, 이들은 아직 학교에 정식으로 공급되지 않은 현장검토본이다. 3학년 1학기과 4학년 1학기 교과서에서 각각 7편의 정보전달 글을 분석 대상으로 선정하였다. 분석 대상의 개요는 다음 <표 1>과 같다.

<표 1> 분석 대상

교과서	제재	교과서	제재
3-1 (가)	옛날에는 어떤 과자를 먹었을까요	4-1 (가)	동물이 내는 소리
	민화		웅기
	플랑크톤이란?		지구가 아프대요
3-1 (나)	신비로운 바닷속 이야기	4-1 (나)	미래의 교통수단
	쓰레기를 줄이자		화성 탐사의 현재와 미래
	사이좋게 지내자		동물 속에 인간이 보여요
	반딧불이		한글이 위대한 이유
합계	7편		7편

2. 분석 도구 및 방법

글의 응집성을 측정하기 위해 통계프로그램 언어 R을 사용하였다. R은 공개용 소프트웨어이다. R로 글을 분석하기 위해 몇 가지 패키지를 사용하였다. KoNLP, tm, lsa 패키지 등을 사용하였다. KoNLP는 형태소 분석을 하기 위해 사용한 패키지이다. KoNLP는 카이스트에서 개발한 한나눔형태소 분석기를 R에서 사용할 수 있도록 연동한 것이다. tm은 형태소로 분석한 글을 행렬로 만들기 위해 사용한 패키지이다. lsa는 SVD를 수행하고, 코사인 유사도를 구하기 위해 사용한 패키지이다.

II장에서 기술한 LSA 수행 단계에 따라 분석을 수행하였다. 먼저, ‘단어 × 문단’의 행렬을 만들었다. 이를 위해 14편의 분석 대상 글을 포함하여 본 연

구에서 사용할 말뭉치를 구축하였다. 말뭉치는 7차 및 2009개정 시기의 초·중고등학교 국어 교과서에 실린 정보전달 글로 이루어졌다. 말뭉치의 크기는 약 6만 어절이었다. 말뭉치를 형태소 단위로 분석한 후에 일반명사, 동사, 형용사를 추출하였다. 그리고 매우 낮은 빈도로 출현한 단어는 분석할 필요가 없으므로, 말뭉치에서 5회 이상 출현한 단어만을 분석 대상으로 하였다. 결과적으로, '1,736단어 × 1,679문단'의 행렬이 만들어졌다. 그 후, 행렬에 가중치를 부과하였는데, 단어 빈도를 로그로 변환한 가중치와 역문서빈도 가중치를 사용하였다.

다음으로, 행렬에 대하여 SVD를 수행하였다. SVD를 수행하는 데 몇 차원으로 축소할지 결정하는 것이 중요하다. Bradford(2008: 155)는 49개의 LSA 관련 선행연구에서 차원의 수를 살펴보았다. 그 결과, 연구물에 따라 최적의 차원 수는 6부터 1,000이상까지를 범위로 하였다. 몇 차원으로 축소할 것인가의 문제는 연구 목적과 말뭉치의 규모 및 특성에 따라 달라진다(Evangelopoulos et al., 2012: 31). 본 연구는 보다 추상적인 수준에서 의미의 핵심을 추출하는 데 목적이 있다. 또한 말뭉치의 규모는 6만 어절로 작은 크기이다. 이러한 점을 고려하여 50개의 차원으로 축소하였다.

끝으로, 코사인 유사도를 활용하여 50차원의 의미 공간에서 분석 대상 글(14편)의 응집성을 측정하였다. 응집성은 '미시 응집성(local coherence)'과 '거시 응집성(global coherence)'으로 나눌 수 있다. 미시 응집성은 주로 문장 단위에서 측정된다. 즉, 인접한 문장 쌍의 유사도를 모두 구한 후에, 이 값들의 평균을 구함으로써 측정된다. 이에 비해 거시 응집성은 주로 문단 단위에서 측정된다. 즉, 인접한 문단 쌍의 유사도를 모두 구한 후에, 이 값들의 평균을 구함으로써 측정된다. 본 연구에서는 거시 응집성을 측정하였다. 예를 들어, 3문단으로 된 글이 있을 때, 첫 번째 문단과 두 번째 문단의 유사도를 구하고 두 번째 문단과 세 번째 문단의 유사도를 구한 후에, 이 값들의 평균을 계산하였다.

IV. 결과 및 논의

초등학교 3~4학년군 국어 교과서에 수록된 정보전달 글의 응집성을 측정한 결과는 다음 <표 2>와 같다.

<표 2> 응집성 측정 결과

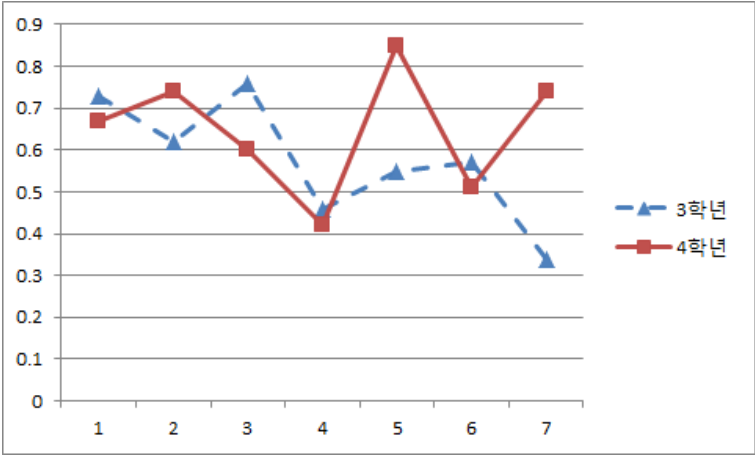
3학년 1학기	응집성	4학년 1학기	응집성
옛날 과거	0.73	동물 소리	0.67
민화	0.62	옹기	0.74
플랑크톤	0.76	지구	0.60
바닷속	0.46	미래 교통	0.42
쓰레기	0.55	화성 탐사	0.85
사이좋게	0.57	동물 속 인간	0.51
반딧불이	0.34	한글	0.74

위의 표에서 볼 수 있는 것처럼, 응집성은 ‘0.34~0.85’를 범위로 하였다. 코사인 유사도로 측정된 응집성은 이론적으로 ‘-1 ~ +1’을 범위로 한다. 하지만 텍스트 분석에서 0 이하의 값을 나타내는 경우는 거의 없다. 따라서 유사도가 1에 가까울수록 응집성이 높은 글이며, 유사도가 0에 가까울수록 응집성이 낮은 글이다. 가장 낮은 응집성을 지닌 글은 3학년 1학기 교과서에 실린 ‘반딧불이’라는 글이다. 낮은 응집성 값은 이 글의 문단별 내용에 큰 격차가 있다는 것을 뜻한다. 응집성 있는 글은 인접 문단 간에 큰 변화나 이동이 적은 ‘여유 있는 산책’이어야 한다(Foltz, 2007: 182). 특히 학생들이 학습 상황에서 새롭게 배우는 글은 응집성이 높을수록 학습가능성이 높다. 가장 높은 응집성을 지닌 글은 4학년 1학기 교과서에 실린 ‘화성 탐사’라는 글이다.

교육 내용의 연계성 측면에서 학년별 평균을 비교해 보는 것도 필요하다. 3학년 1학기 글 7편의 평균 응집성은 0.58이었고, 4학년 1학기 글 7편의

평균 응집성은 0.65였다. 학년이 올라감에 따라, 글의 응집성이 점차 낮아지는 것이 이상적이다. 하지만 본 연구에서 분석한 결과는 ‘후퇴(낙차) 현상’을 나타냈다. 후퇴(낙차) 현상은 직전 학년의 교과서 내용이 직후 학년의 교과서 내용보다 어려울 때 생긴다(천경록, 2016: 241). 응집성 측면에서만 보았을 때, 3학년 교과서가 4학년 교과서보다 더 어려웠다.

한 학년 내에서도 단원이 전개됨에 따라, 글의 응집성이 점차 낮아지는 것이 이상적이다. 이를 확인하기 위해 앞에서 제시한 <표 2>를 그래프로 나타내면, 다음 <그림 4>와 같다.



<그림 4> 학년별 응집성 측정 결과

위의 그림에서 가로축은 학년별로 7개의 제재가 제시되는 순서이다. 세로축은 코사인 유사도이다. 4학년에 비하여, 3학년(점선)은 학습 내용이 전개됨에 따라 응집성이 점차 낮아지는 경향을 보인다. 반면에, 4학년은 마지막에 제시되는 제재의 응집성(0.74)이 처음에 제시되는 제재의 응집성(0.67)보다 더 높다. 제재 선정 과정에서 응집성에 대한 고려가 필요해 보인다.

V. 결론

본 연구의 목적은 글의 수준을 평가하기 위해 글의 응집성을 측정하는 데 있다. 연구 결과를 요약하면 다음과 같다. 학생들에게 적정 수준의 글을 제공하는 것은 읽기 교육 분야의 오랜 관심사이다. 이는 어떠한 언어적 변수가 글을 어렵게 하는지 또는 쉽게 하는지 찾는 것으로부터 시작한다. 그런데 심층 수준의 변수인 글의 응집성은 측정의 어려움으로 인하여 소홀히 다루어져 왔다. 응집성에 대한 평가는 주로 교사나 전문가의 질적 평가에 맡겨져 왔다. 하지만 질적 평가는 주관적일 수밖에 없다는 단점을 지닌다. 질적 평가는 평가자들 사이에 큰 편차를 설명하지 못한다. 이를 대체하는 양적 평가 방법은 명제 분석을 활용하는 것이었다. 하지만 이 방법은 글을 명제 단위로 분석하는 과정에서 오랜 시간과 수고를 요구한다.

그래서 본 연구에서는 응집성을 측정하는 데 기존 양적 평가의 대안적인 방법으로 잠재의미분석(LSA)에 주목하였다. LSA의 특성 및 기본 가정을 살펴본 후에, 초등학교 교과서에 수록된 글을 시험적으로 분석해 보았다. 분석 결과는 응집성 측면에서 제재에 대한 재검토가 필요함을 나타냈다.

LSA를 활용한 응집성의 측정은 국어교육 분야에서는 아직 낯설다. 하지만 영어권에서는 이 방법이 정착되어 왔다. 예를 들어, 뎀피스 대학교에서 개발한 Coh-Metrix 시스템은 영어로 된 글에 대한 여러 가지 지수를 산출하는데, 그중 하나로 코사인 유사도를 제시한다. 국내의 영어 교육 분야에서도 이와 관련된 연구들을 살펴볼 수 있다(예, 김정렬·양지윤, 2012; 전문기, 2011; 전문기·임인재, 2009). Coh-Metrix 시스템에 접근하여 영어 글을 입력하면 결과를 쉽게 얻을 수 있기 때문에, 국내 영어 교과서의 난이도는 응집성 측면에서도 분석되고 있다. 앞으로 국어교육 분야에서도 응집성 측정에 대한 논의가 활발히 이루어지기를 기대한다.

* 본 논문은 2017. 2. 1. 투고되었으며, 2017. 2. 14. 심사가 시작되어 2017. 3. 9. 심사가 종료되었음.

참고문헌

- 강범모(2014), 「텍스트 맥락과 단어 의미: 잠재의미분석」, 『언어학』 62, 3-34, 한국언어학회.
- 김명순(2010), 「읽기 교육과정에 나타난 텍스트성 고려 양상」, 『독서연구』 24, 179-208, 한국독서학회.
- 김정렬·양지윤(2012), 「Coh-Metrix를 통한 초·중등 영어교과서 연계성 분석」, 『영어교육』 62(2), 319-341, 한국영어교육학회.
- 이경남(2017), 「잠재적 의미 분석(LSA)를 활용한 독해 과정에서 추론 양상 분석」, 『독서연구』 42, 107-131, 한국독서학회.
- 이재승(2003), 「읽기와 쓰기 행위에서 결속 구조의 의미와 지도」, 『국어교육』 110, 91-111, 한국국어교육학회.
- 전문기(2011), 「Coh-Metrix를 이용한 중학교 1학년과 2학년 개정 영어교과서 읽기 자료의 코퍼스 언어학적 연계성 분석」, 『언어과학연구』 56, 201-218, 언어과학회.
- 전문기·임인재(2009), 「코메트릭스(Coh-metrix)를 이용한 중학교 1학년 개정 영어 교과서의 코퍼스 언어학적 비교 분석」, 『영어교육연구』 21(4), 265-292, 팬코리아영어교육학회.
- 천경록(2016), 「초등학교와 중학교 국어 교과서 간의 접합성 분석」, 『새국어교육』 107, 237-265, 한국국어교육학회.
- Bradford, R. B.(2008), "An empirical study of required dimensionality for large-scale Latent Semantic Indexing applications," *CIKM-INTERNATIONAL CONFERENCE* - 17(1), 153-162.
- Evangelopoulos, N. et al.(2012), "Latent semantic analysis: five methodological recommendations," *European Journal of Information System* 21(1), 1-45.
- Foltz, P. W.(1996), "Latent semantic analysis for text-based research, Behavior Research Methods," *Instrument, & Computers* 28(2), 197-202.
- Foltz, P. W.(2007), "Discourse coherence and LSA," In: T. K. Landauer, D. S. McNamara, S. Dennis, & W. Kintsch (eds.), *Handbook of Latent Semantic Analysis* (pp. 167-184), Mahwah, NJ: Lawrence Erlbaum Associates, Inc.
- Foltz, P. W. et al.(1998), "The measurement of textual coherence with Latent Semantic Analysis," *Discourse Processes* 25, 285-307.
- Kintsch, W. et al.(1975), "Comprehension and recall of text as a function of content variables," *Journal of Verbal Learning and Verbal Behavior* 14(2), 196-214.
- Landauer, T. K.(2007), "LSA as a theory of meaning," In: T. K. Landauer, D. S. McNamara, S. Dennis, & W. Kintsch (eds.), *Handbook of Latent Semantic Analysis* (pp. 3-34), Mahwah, NJ: Lawrence Erlbaum Associates, Inc.
- Landauer, T. K. et al.(1998), "An introduction to Latent Semantic Analysis," *Discourse Processes* 25, 259-284.

- Martin, D. I. & Berry, M. W. (2007), "Mathematical foundations behind latent semantic analysis," In: T. K. Landauer, D. S. McNamara, S. Dennis, & W. Kintsch (eds.), *Handbook of Latent Semantic Analysis* (pp. 35-55), Mahwah, NJ: Lawrence Erlbaum Associates, Inc.
- McNamara, D. S. et al. (2010), "Coh-Metrix: Capturing linguistic features of cohesion," *Discourse Processes* 47 (4), 292-330.
- Meyer, B. J. F. (1982), "Reading research and the composition Teacher: The importance of plans," *College Composition and Communication* 33(1), 37-49.
- Quesada, J. (2007), "Creating your own LSA space," In: T. K. Landauer, D. S. McNamara, S. Dennis, & W. Kintsch (eds.), *Handbook of Latent Semantic Analysis* (pp. 71-85), Mahwah, NJ: Lawrence Erlbaum Associates, Inc.
- Todirascu, A. et al. (2013), "Coherence and cohesion for the assessment of text readability," *Natural Language Processing and Cognitive Science* 11, 11-19.

글의 수준을 평가하기 위한 글의 응집성 측정

조용구

본 연구의 목적은 글의 수준을 평가하기 위해 글의 응집성을 측정하는 데 있다. 응집성은 글의 수준에 영향을 주는 심층의 변수이다. 하지만 측정의 어려움으로 인하여 글의 수준 평가에서 소홀히 다루어져 왔다. 이에 따라, 글의 응집성을 측정하는 방법으로서 잠재의미분석(LSA)에 대하여 기술하였다. 그 후, 초등학교 국어 교과서에 수록된 정보전달 글의 응집성을 측정해보았다. 분석 결과는 교과서에 수록된 글의 수준이 응집성의 측면에서 재검토될 필요가 있음을 나타냈다.

핵심어 글의 수준, 글의 응집성, 잠재의미분석(LSA), 국어 교과서, 정보전달 글

ABSTRACT

The measurement of Text Coherence to Evaluate the Text Difficulty

Jo Yonggu

The purpose of this study is to measure the text coherence in order to evaluate the text difficulty. The text coherence is deeper level of language variable that have a influence on the text difficulty. But it has been neglected in the assessment of the readability due to the difficulty of measurement. Therefore, the study introduced Latent Semantic Analysis (LSA) as a alternative technique to measure the coherence. And then, 14 Informational texts of Korean Language Textbook were analyzed. The results showed informational texts of Korean Language Textbook were needed to revise on the text difficulty.

KEYWORDS Text Difficulty, Text Coherence, Latent Semantic Analysis, Korean Language Textbook, Informational text