

독자에게 적합한 텍스트 선정 방법의 개발 방향

조용구 광주하남중앙초등학교 교사

- I. 서론
- II. 이론적 배경
- III. 주요 과제와 개발의 방향
- IV. 결론

I. 서론

읽기 교육에서 독자의 수준에 적합한 텍스트를 제공하는 일은 읽기 교육 연구자나 교사들의 주요 관심사이다. 그러나 독자의 읽기 능력을 정확하게 측정하는 것은 물론 텍스트의 난이도를 측정하는 것의 어려움으로 인해, 독자에게 맞는 텍스트를 제공하는(matching) 일은 쉽게 이루어지지 않고 있다. 즉, 독자에게 적합한 텍스트를 제공하는 일은 주로 교사나 읽기 교육 전문가의 경험이나 직관에만 의존하고 있는 실정이다.

이와 관련하여 책을 고르는 것을 신발을 고르는 것에 비유한 다음의 두 사례를 살펴보자. 먼저, Anderson & Armbruster(1986)에 제시된 사례이다. 발의 크기는 자로 잴 수 있으며 신발 역시 정확한 크기로 만들어지기 때문에, 발과 신발의 크기만 알면 신발을 쉽게 살 수 있다. 그러나 학생의 읽기 능력은 정확하게 측정하기 어려우며, 텍스트의 난이도도 신발의 크기만큼 정확하게 판단할 수 없다. 따라서 학생의 읽기 능력 평가와 이독성 공식에 의한 텍스트의 난이도 측정은 믿을 만하지 못하다.

다음으로, Schnick & Knickelbine(2000: 6-7)에 제시된 사례이다. 한

사람이 아들의 신발을 사기 위해, 가게에 가서 점원에게 5학년 신발을 달라고 했다. 점원은 수상한 눈빛으로 쳐다보며 말했다. 5학년 학생들의 발의 크기가 다양하기 때문에 5학년 신발을 줄 수 없다는 것이다. 또한 그 가게는 발의 크기와 신발의 크기를 정확한 수치로 재서 신발을 권하고, 이것은 기본이기 때문에 더 중요한 것은 신발의 성능, 스타일, 색상 등이라는 것이다. 점원의 말을 듣고, 그는 크기가 8이고 발목이 높은 빨간색 농구화를 구입했다. 잠시 후, 그는 신발 가게 옆에 있는 서점으로 갔다. 서점에서 5학년 학생에게 맞는 책을 추천해 달라고 했다. 서점 직원은 망설임 없이 ‘5학년 권장도서’ 칸에 가서 책 한 권을 추천해 주었다. 아들에게 농구화와 책 한 권을 선물로 주었는데, 아들은 책이 어려워 읽지 않고 농구 연습에만 열중하였다.

위의 첫 번째 사례는 학생의 읽기 능력과 텍스트의 난이도를 측정하기 어렵다는 입장이다. 학생의 읽기 능력은 배경지식, 동기, 흥미, 읽기 기능이나 전략 등 다양한 요인들로 이루어져 있으며, 텍스트의 난이도 또한 기존의 이독성 공식에서 다룬 단어의 빈도, 문장의 길이 외에도 여러 가지 요인들이 복합적으로 작용한다는 것이다. 따라서 이러한 입장에서는 텍스트를 선정할 때 전문가의 질적 판단을 중요시한다. 그러나 전문가의 경험이나 직관에만 의존하는 것은 주관적이라는 비판을 피하기 어렵다.

두 번째 사례는 독자에게 적합한 텍스트를 선정하는 일이 어렵지만, 발과 신발의 크기를 재는 것과 같이 독자와 텍스트를 정확한 수치로 재는 일이 반드시 필요하다고 주장한다. 신발을 구입할 때, 발과 신발의 크기를 재는 것은 중요한 일이며 이미 기본적인 사항이 되어 있다. 그러나 책을 선정하는 것은 독자와 텍스트를 측정하는 것이 어렵다는 이유만으로 소홀히 다루어져 왔다.

따라서 이 논문의 목적은 독자의 읽기 능력과 텍스트의 난이도를 정확하게 측정하여, 독자에게 적합한 텍스트를 선정하는 방법의 개발 방향을 모색하는 데 있다. 이를 위해 II장에서는 독자의 읽기 능력과 텍스트 난이도의 측정에 관련된 선행연구들을 검토한다. III장에서는 독자와 텍스트를 측정하

는 데 관련된 주요 과제와 그에 따른 해결 방안을 모색한다.

II. 이론적 배경

이 장에서는 독자의 읽기 능력¹과 텍스트의 난이도를 측정하는 것에 관련된 국내외의 선행연구들을 검토한다. 이 중 독자와 텍스트를 측정하기 위해, 본고에서 특히 관심을 갖는 Lexile 지수는 절을 달리하여 살펴본다.

1. 선행연구 검토

먼저, 국내의 선행연구들을 살펴보기로 한다. 이성영(2011)은 초등학교 국어, 사회, 과학 교과서의 이독성(readability)을 비교하였다. 이독성은 ‘글을 읽고 이해하기 쉬운 정도’를 나타내는 개념이다. 그는 단어의 빈도와 문장의 길이를 바탕으로 각 교과서의 이독성을 측정하였다. 이 때, 단어의 빈도를 측정하기 위하여 김광해(2003)의 교육용 어휘 등급을 활용하였다. 이 어휘 등급은 어휘를 1~4등급으로 분류한 것이다. 또한 문장의 길이를 측정하기 위하여 ‘문장당 평균 어절 수’를 계산하였다. 이는 전통적인 이독성 공식에서 가장 빈번하게 활용하는 방법이기도 하다.

박순원(2011)은 (주)날말에서 개발한 ‘독서지수’를 활용하여, 텍스트의 난이도를 측정하였다. 이 독서지수 프로그램은 약 50만 단어를 1~9등급으

1 읽기 능력(reading ability)은 해독 기능(decoding skills)과 독해 전략(comprehension strategies)을 포괄하는 개념이다. 그런데 우리나라의 경우, 취학 전에 해독 기능을 이미 습득한 학생들이 많기 때문에, 여기에서의 읽기 능력은 주로 독해 전략을 의미한다. 기능, 전략, 능력의 개념에 대한 자세한 내용은 천경록(1995, 2009), Afflerbach, Pearson & Paris(2008)를 참고.

로 분류한 데이터베이스를 가지고 있다. 예를 들어, ‘남자’, ‘선생님’과 같은 단어는 1등급의 쉬운 단어이고, ‘군납하다’, ‘해뜨리지다’와 같은 단어는 9등급의 어려운 단어이다. 이 프로그램에 기초하여, 텍스트는 100~1,850 범위의 독서지수를 부여받고, 학년별·단계별 권장도서를 선정할 수 있게 된다. 이 방법은 대량의 코퍼스를 활용하며 프로그램이 자동화되어 있다는 장점이 있다. 그러나 텍스트의 난이도를 측정하는 데 지나치게 단어의 빈도에만 의존한다는 문제점을 가지고 있다.

김대희(2011)는 독자의 수준과 텍스트의 수준을 위계화하는 방안을 비판적으로 고찰하였다. 검사 도구에 의해 독자의 수준을 구분하는 방법으로 국제 학업성취도 평가 프로그램(PISA), 미국의 국가수준 학업성취도 평가(NAEP) 등을 검토하였다. 그는 이러한 방법들이 상대적인 등급화에 관심이 있기 때문에, 독자의 수준을 절대적인 척도로 나타내지 못하는 문제점이 있다고 하였다. 검사 도구에 의해 텍스트의 수준을 위계화하는 방법으로는 Lexile 지수와 (주)날말의 독서지수를 검토하였다. 그리고 텍스트의 수준을 나누는 데 텍스트의 다양한 변인을 고려하지 못한다는 점, 텍스트의 난이도 지수가 10단위로 너무 세분화되어 있다는 점 등을 문제점으로 지적하였다.

최숙기(2012)는 읽기 교육용 텍스트를 선정하는 방법을 고안하였다. 텍스트의 난이도를 양적으로 측정하기 위하여, 전통적 이독성 공식들을 검토한 후, 저학년(1~5학년)과 고학년(6~10학년)에 따라 서로 다른 이독성 공식을 활용하는 것이 적절하다고 하였다. 이 연구는 여러 가지 이독성 공식을 실제 텍스트에 적용하였다는 데에 의미가 있다. 그러나 이독성 공식은 텍스트 자체의 특성만을 고려한다는 비판을 받아왔다. 우리나라 학생들에게 적합한 이독성 공식을 제안하는 것도 독자의 수준을 고려하지 못한다는 문제점을 안고 있다.

다음으로, 국외의 선행연구 몇 가지를 살펴보기로 한다. 국외의 선행연구 중 주목할 만한 것은 2010년에 발표된 미국의 국가수준 교육과정인 Common Core State Standards(CCSS; 이하 중핵성취기준)이다. 중핵성취

기준에서 읽기 영역을 살펴보면, 학년별로 10개의 성취기준을 제시하고 있는데, 이 중 1~3번 성취기준은 ‘중심 생각과 세부 내용’, 4~5번은 ‘기교와 구조’, 7~9번은 ‘지식과 생각의 통합’, 10번은 ‘읽기의 범위와 텍스트 복잡도의 수준’이다. 읽기의 범위와 텍스트 복잡도를 하나의 성취기준으로 제시하고 있다. 2학년과 3학년의 10번 성취기준을 예로 들면, 다음 <표 1>과 같다.

표 1. 읽기 영역 10번 성취기준(2~3학년 정보전달 텍스트)

2학년	2학년말까지, 2·3학년 텍스트 복잡도군에서 능숙하게, 가장 높은 범위에서 필요하다면 비계를 제공받으며, 역사/사회, 과학, 기술 텍스트를 포함한 정보전달 텍스트를 읽고 이해한다.
3학년	3학년말까지, 2·3학년 텍스트 복잡도군의 최상에서 독립적으로 그리고 능숙하게, 역사/사회, 과학, 기술 텍스트를 포함한 정보전달 텍스트를 읽고 이해한다.

2학년의 경우, 2학년말까지 읽기의 범위는 역사/사회, 과학, 기술 텍스트를 포함한 정보전달 텍스트이다. 텍스트 복잡도의 수준은 2~3학년 텍스트 복잡도군에 제시된 텍스트를 능숙하게 읽으며, 이 중 수준이 높은 텍스트는 교사의 도움을 받으며 읽도록 하고 있다. 3학년의 경우, 읽기의 범위는 2학년과 동일하다. 텍스트 복잡도의 수준은 2~3학년 텍스트 복잡도군에 제시된 텍스트를 독립적이며 능숙하게 읽도록 하고 있다. 이러한 성취기준과 관련하여 <부록 B>에서는 2~3학년 텍스트 복잡도군의 정보전달 텍스트 21편을 예로 제시하고 있다.

국외의 연구 중 텍스트 난이도와 관련한 다른 연구로 Coh-Metrix가 있다. Coh-Metrix는 멤피스 대학에서 운영하는 비영리 서비스로, 텍스트를 입력하면 60가지 이상의 텍스트 관련 지수를 얻을 수 있다. Coh-Metrix는 단어 수준, 통사론, 명제, 공지시어와 결속성(coherence), 갈래에 관한 지수들을 제공한다(Graesser, McNamara, Louwerse, 2011: 44-48). 이러한 분석 방법은 자동화되어 있을 뿐만 아니라, 전통적인 이독성 공식에서 다루지 못했던 텍스트의 다양한 측면을 측정할 수 있으므로 큰 장점을 가지고 있다. 그러나 각각의 지수가 복잡하고, 이 지수들은 같은 텍스트에 대하여 상반

된 분석 결과를 내놓기도 하기 때문에 활용도가 떨어진다. 이에 따라 Coh-Metrix 개발자들은 텍스트의 핵심적인 특성을 드러내는 요인을 분리해 내는 연구를 지속적으로 수행해 오고 있지만, 아직 일반인들이 활용하기에는 어려움이 있으며, 따라서 중핵성취기준에서도 이를 반영하지 않고 있다(CC-SSI, 2010: 7). 현재 Coh-Metrix는 텍스트의 다양한 측면을 고려하여 자동화된 방법으로 텍스트의 난이도를 측정할 수 있는 가능성을 보여 주었다는 점에서 의의를 찾을 수 있다.

지금까지 독자의 읽기 능력과 텍스트 난이도의 측정에 관련된 국내외의 선행연구 몇 가지를 살펴보았다. 최근 이루어진 국내의 연구들은 독자와 텍스트 측정에 관한 기존의 연구들을 비판적으로 고찰하거나, 이독성 공식에 기초하여 텍스트의 난이도를 측정하는 것들이다. 국외의 연구로 미국의 중핵성취기준은 텍스트 복잡도를 성취기준의 하나로 제시하고 있으며, 그 이론적 근거로 Lexile 지수를 활용한다. 또한 아직 완벽하지는 않지만 Coh-Metrix와 같이 텍스트의 다양한 층위를 분석하는 시도도 이어지고 있다. 다음 절에서는 Lexile 지수에 대하여 살펴본다.

2. Lexile 지수

Lexile는 MetaMetrics(2000)에서 개발한 것으로, 학생의 읽기 능력에 따라 적절한 난이도의 텍스트를 제공하기 위한 프로그램이다. 이 프로그램은 독해 검사를 실시하여 학생의 읽기 능력을 파악하고 Lexile 분석기로 텍스트의 난이도를 측정한 후에, 독자에게 적합한 텍스트를 제공하는 것을 목표로 한다. Lexile의 기본 가정은 다음 <그림 1>과 같다.

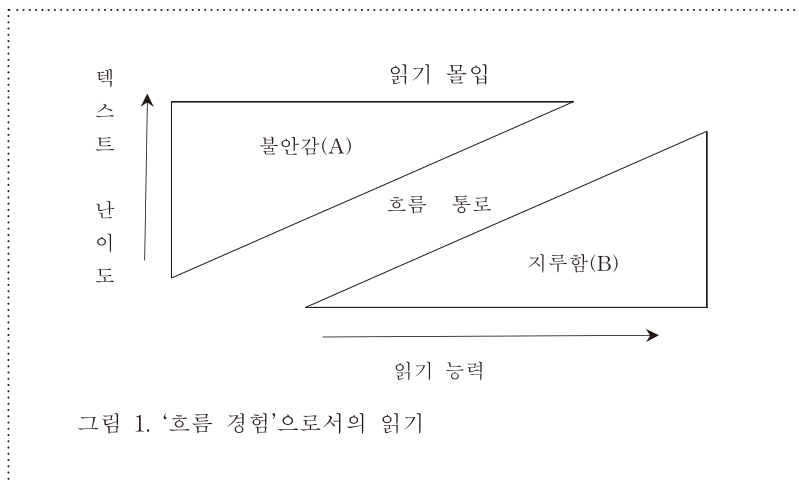


그림 1. '흐름 경험'으로서의 읽기

위의 그림에서는 읽기를 독자가 자신감과 성취감을 느끼는 '흐름 경험'이라고 표현한다(Schnick & Knickelbine, 2000: 4-6). 몰입하여 글을 읽는 것이 흐름을 가진다는 것이다. 자신의 읽기 능력에 비해 너무 어려운 텍스트를 접하게 되면, 불안감을 느낀다(A부분). 반면에 자신의 읽기 능력에 비해 너무 쉬운 텍스트를 접하게 되면, 지루함을 느낀다(B부분). 읽기에 몰입할 수 있는 것은 자신의 능력에 적절한 수준의 텍스트가 제공되었을 때이다(흐름 통로). 읽기에서 '흐름' 상태를 만드는 것은 학생의 읽기 능력에 적절한 텍스트를 제공하는 것이다.

독자에게 적합한 난이도의 텍스트를 제공하는 것의 중요성은 지도적 수준의 개념과 관련된다(Walpole, Hayes & Robnolt, 2006: 2). 학생과 텍스트의 관계에 따라, 읽기의 수준을 독립적 수준(independent level), 지도적 수준(instructional level), 좌절 수준(frustration level)으로 나눌 수 있다.² 어떤

2 이를 학년 상회 수준(above grade), 학년 수준(at grade), 학년 미달 수준(below grade)이라고도 한다.

학생에게 어떤 텍스트는 너무 쉽고(독립적 수준), 어떤 텍스트는 교사의 도움을 필요로 하며(지도적 수준), 어떤 텍스트는 혼자 읽기에 너무 어렵다(좌절 수준). 위의 그림에서 A부분은 좌절 수준이며, B부분은 독립적 수준이다. 독자에게 유의미한 읽기 지도(흐름 통로)가 이루어지기 위해서는 지도적 수준의 텍스트가 제공되어야 한다.

Lexile는 차세대(new generation) 이독성 시스템이라고 불린다(Hiebert, 2012: 14; 2013: 461). 차세대 이독성 시스템이라고 부르는 것은 기존의 이독성 공식이 지닌 문제점을 조금씩 극복하고 있기 때문이다. Lexile가 기존의 이독성 공식과 다른 점은 다음과 같다.

첫째, 텍스트 요인과 함께 독자 요인을 중요하게 고려한다. 읽기가 독자와 텍스트의 상호작용 혹은 교섭의 과정임에도 불구하고, 기존의 이독성 공식들은 텍스트 요인만을 평가한다는 비판을 받아왔다. 이에 반해 Lexile는 독자 요인과 텍스트 요인을 두 축으로 설정함으로써, 진일보한 모습을 보인다.

둘째, 개별 학생의 읽기 능력과 텍스트의 난이도를 동일한 척도에 놓는다. 즉, 학생의 읽기 능력과 텍스트의 난이도를 0~2,000까지의 Lexile(L) 지수로 제시한다. 이를 그림으로 나타내면 다음 <그림 2>와 같다.

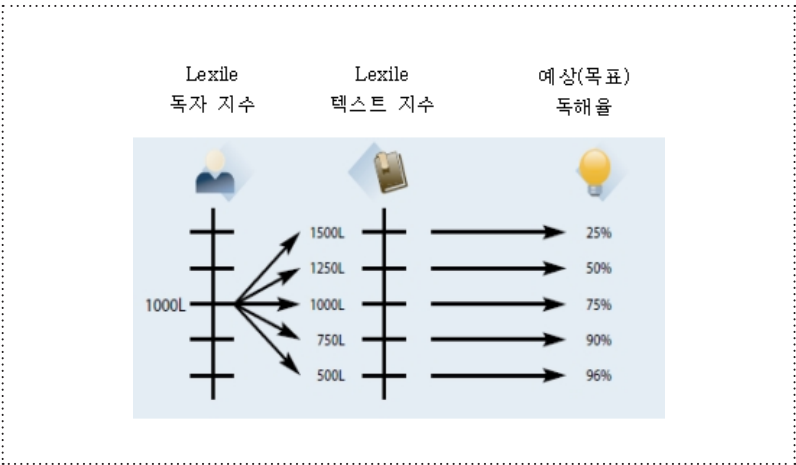


그림 2. Lexile 독자 지수와 텍스트 지수

위의 <그림 2>에서 볼 수 있는 것처럼, 읽기 능력이 1,000L인 학생은 텍스트 난이도가 1,000L인 책을 읽고 75% 정도를 이해할 수 있는 것으로 예상된다. 이 학생에게 1,500L의 책은 너무 어려워서 25% 정도밖에 이해가 되지 않고, 500L의 책은 너무 쉬워서 96% 정도를 이해할 수 있다. 75%의 독해율은 목표 독해율로서, 독립적 읽기를 수행할 때 목표로 하는 비율이다. 만약 이 학생이 교사의 도움을 받아 책을 읽는다면, 독해율은 좀 더 증가할 것이다. 이 목표 독해율 75%는 독자가 텍스트를 충분히 이해할 수 있으면서, 동시에 약간의 도전감을 느낄 수 있는 지점이다. 다시 말해, 독자는 너무 쉬운 텍스트로 인한 지루함을 느끼지 않고, 너무 어려운 텍스트로 인한 불안감을 경험하지 않는다. 그 결과, 가치 있는 읽기를 경험할 수 있다.

셋째, 학생과 텍스트를 동일한 척도에 놓음으로써, 교육적 의사 결정을 하기가 쉬워진다. Lexile는 학생의 읽기 능력에서 100이하 50이상의 Lexile 지수를 나타내는 텍스트를 권장한다. 예를 들어, 어떤 학생의 읽기 능력이 1,000L이라면, 900L~1,050L의 텍스트를 제공한다. 여기에서 두 가지 시사점을 얻을 수 있다. 먼저, 학년 수준(at grade)의 독자에게 적절한 텍스트를 제공할 수 있다. 예를 들어, 우리나라 3학년 학생들의 평균 읽기 능력이 500L이라면, 교과서에 제시되는 텍스트의 수준을 400~550L로 정할 수 있다. 다음으로, 동일 학년 내에 존재하는 다양한 수준의 독자에게 적절한 난이도의 텍스트를 제공할 수도 있다. 우수한 독자에게는 좀 더 어려운 텍스트를, 부진한 독자에게는 좀 더 쉬운 텍스트를 제공할 수 있는 것이다.

끝으로, Lexile는 어휘에 관한 정보를 얻기 위해 디지털 기술의 장점을 활용한다(Hiebert, 2013: 461). 거의 모든 이독성 공식들은 텍스트의 2가지 주요 특성을 분석한다. 문장과 어휘가 그것이다. 문장을 계산할 때, 비록 소수의 공식들이 음절수를 계산하기는 하지만 거의 대부분은 ‘문장 당 평균 단어 수’를 계산한다. 어휘에 관하여, 어떤 공식들(예, Spache, 1953)은 텍스트의 단어들을 학년별 어휘 목록과 비교하는 반면, 다른 공식들(예, Fry, 1968)은 ‘단어 당 평균 음절수’를 복잡도의 지표로 사용한다. 이에 반해, Lexile는 ‘단

어의 빈도수'를 계산하는데, 이는 해당 텍스트에서 단어가 출현하는 빈도가 아니라, 광범위한 코퍼스의 데이터베이스와 비교하여 계산되는 빈도수이다.

이처럼 Lexile는 기존의 이독성 공식에 비하여 진일보한 모습을 보인다. 독자 요인을 중요하게 고려하고 있으며, 독자와 텍스트를 동일한 척도에 놓음으로써 교육적 의사결정의 가능성을 높이고 있다.

III. 주요 과제와 개발의 방향

이 장에서는 독자에게 적합한 텍스트 선정 방법을 개발하기 위하여, 고려해야 할 주요 과제와 그에 따른 개발 방향을 모색한다. 여기에서는 Lexile 개발 최종보고서인 Smith, Stenner, Horabin & Smith(1989)를 참고한다. Lexile와 같은 방법을 개발하기 위해서는 문항 개발과 문항난이도 분석, 회귀분석, 척도 변환 등의 과제를 해결해야 한다. 이들을 절을 달리하여 살펴보기로 한다.

1. 문항 개발과 문항난이도 분석

Lexile는 독해 검사 문항으로 학생의 읽기 능력을 측정한다. Lexile와 같은 방법을 개발하기 위한 첫 번째 과제는 문항을 개발하고 문항난이도를 분석하는 것이다.

먼저, 독해 검사 문항을 개발하여야 한다. Lexile에서 사용하는 독해 검사 문항은 보통 4지 선다형으로 제시되며, 문장에서 빠진 단어를 채우는 빈칸메우기 검사이다(Graesser *et al.*, 2011: 44; Walpole *et al.*, 2006: 6). 검사 문항 하나를 예로 들면, 다음 <그림 3>과 같다(Smith *et al.*, 1989: 11).

월버는 매일 점점 더 잘못을 좋아했다. 곤충에 대항하는 그녀의 캠페인은 합리적이고 유용해 보였다. 농장 근처의 누구도 파리에게 좋은 말을 해 줄 수 없었다. 파리들은 다른 사람들을 난처하게 하는 데 시간을 소비했다. 소들은 그들을 경멸했다. 말들도 그들을 경멸했다. 양들은 그들을 몹시 싫어했다. 주커만씨 부부는 항상 그들에게 불평하고, 비명을 지르고 있었다. 모든 사람들이 그들에게 _____.

- a. 동의했다 b. 찌푸렸다 c. 비웃었다 d. 배웠다.

그림 3. Lexile 검사 문항의 예시

다음으로, 개발한 문항에 대하여 문항난이도(difficulty)³를 분석한다. 문항난이도는 문항의 어려운 정도를 나타내는 지수이다. 독해 검사는 문항난이도로 학생들 간 읽기 능력의 차이를 밝힐 수 있다. 쉬운 문항만 맞힌 학생은 읽기 능력이 낮으며, 어려운 문항까지 맞힌 학생은 읽기 능력이 높다.

검사 문항의 분석과 관련하여, 읽기 능력 표준화 검사 도구를 떠올려 볼 수 있다. 이러한 표준화 검사 도구에는 천경록(2006), 한철우·이경화·최규홍(2007), 윤준채·서혁(2010) 등이 있다. 이 검사 도구들은 독자의 읽기 능력을 표준화된 방법으로 측정한다는 점에서 큰 의미가 있다. 그런데 이러한 검사 도구들은 주로 고전검사이론에 기초하여 개발된 것이다. 고전검사이론은 비교적 간단한 절차로 문항을 분석할 수 있는 장점이 있지만, 피험자 능력의 불변성, 문항 모수의 불변성을 극복하지 못하는 한계를 지니고 있다(성태제, 2002: 249-250).

3 최영환(2003: 234)에서 논의한 것처럼, '난이도'라는 용어는 불분명하기 때문에 '난도'라는 용어가 더 적절하다. 그러나 평가 및 측정 분야에서는 난이도라는 용어를 보편적으로 사용하기 때문에, 이 논문에서도 난이도라는 용어를 사용하였다.

이와 같은 고전검사이론의 문제점을 해결하기 위하여 등장한 것이 문항 반응이론이다. 문항반응이론은 문항 하나하나에 근거하여 문항을 분석하고, 피험자의 능력을 추정하는 이론이다. Lexile에서는 문항반응이론에 기초하여 문항을 분석한다. 문항반응이론으로 문항을 분석하면, 문항난이도와 같은 문항모수가 변하지 않기 때문이다. 따라서 본 연구에서도 문항반응이론에 근거하여 문항난이도를 분석하는 것을 개발 방향으로 한다.

문항반응이론에서 문항난이도는 -2부터 +2 사이에 존재하며, 값이 클 수록 문항이 어렵다고 해석한다. 문항난이도에 대한 절대적인 기준은 없으나 대략 다음과 같이 표현한다.

표 2. 문항반응 이론의 문항난이도

문항난이도 지수	언어적 표현
-2.0 미만	매우 쉽다
-2.0 이상 ~ -0.5 미만	쉽다
-0.5 이상 ~ +0.5 미만	중간이다
+0.5 이상 ~ +2.0 미만	어렵다
+2.0 이상	매우 어렵다

문항난이도는 문항반응모형에 따라 분석한다. 일반적으로 사용하는 문항반응모형은 수식이 간단한 로지스틱 모형이다. 로지스틱 모형은 문항난이도만을 고려하는 1 - 모수 로지스틱 모형, 문항난이도와 문항변별도를 고려하는 2 - 모수 로지스틱 모형, 문항난이도와 문항변별도, 문항추측도까지 고려하는 3 - 모수 로지스틱 모형으로 나뉜다.⁴ 본 연구에서는 난이도에 따라 독자와 텍스트를 맞추기 위하여 문항난이도만을 분석하기 때문에, 1 - 모수 로지스틱 모형을 선택한다.

4 문항변별도는 문항이 피험자의 능력 수준을 변별하는 정도를 나타낸다. 문항추측도는 능력이 전혀 없는 피험자가 문항의 답을 추측하여 맞힐 수 있는 정도를 나타낸다.

Lexile는 'Rasch의 문항반응이론'을 사용하는데, Rasch의 문항반응이론은 1-모수 로지스틱 모형과 같다. 1-모수 로지스틱 모형으로 문항난이도를 분석하는 컴퓨터 프로그램으로 RaschAn이 있다. 본 연구에서는 RaschAn 프로그램을 활용하여, 1-모수 로지스틱 모형으로 문항난이도를 분석하는 방향을 선택한다.

정리하면, 독자의 읽기 능력에 적합한 텍스트를 선정하기 위한 첫 번째 과제는 독해 검사 문항을 개발하고, 문항의 난이도를 분석하는 것이다. 문항난이도를 분석하는 데 문항반응이론을 활용한다. 이는 고전검사이론에 비하여 문항난이도를 보다 정확하게 분석할 수 있기 때문이다. RaschAn 프로그램을 활용하여, 1-모수 로지스틱 모형으로 문항난이도를 분석하기로 한다.

2. 회귀분석

두 번째 과제로 회귀분석을 활용하여, 텍스트의 난이도를 측정할 수 있는 회귀방정식을 구해야 한다. 여기에서의 텍스트는 일단 검사 문항에서의 지문을 뜻한다. 검사 문항의 지문에서 구한 회귀방정식은 나중에 다른 텍스트의 난이도를 측정하기 위하여 사용된다.

텍스트의 난이도는 의미적 복잡도와 통사적 복잡도의 두 측면을 고려한다. 의미적 복잡도는 단어의 친숙한 정도에 따른 것이다. 즉, 독자가 일상생활에서 자주 접하게 되는 단어는 더 쉽게 이해할 수 있고, 자주 접하지 못하는 단어는 어렵다는 것이다. 이는 곧 단어의 빈도로 계산된다. 앞서 언급한 대로, 단어의 빈도는 분석 대상 텍스트에서 단어가 나타나는 빈도가 아니라 코퍼스에서의 빈도이다. 통사적 복잡도는 문장인데, 문장의 길이가 짧으면 좀 더 이해하기 쉽고, 문장의 길이가 길면 이해하기 어렵다는 것이다.

단어의 빈도와 문장의 길이로 텍스트의 난이도를 측정하기 위해서 회귀분석을 한다. 회귀분석은 '종속변수와 독립변수들 사이의 선형함수적 관계를 밝히는 통계적 기법'이다(강상진, 2003: 7). 여기에서 종속변수(Y)는 텍

트의 난이도이며, 독립변수(X)는 단어의 빈도와 문장의 길이이다. 독립변수가 1개이면 단순회귀분석이고, 독립변수가 2개 이상이면 중다회귀분석이라고 하는데, 여기에서는 독립변수가 2개이므로 중다회귀분석법을 사용한다.

회귀분석의 함수관계를 설명하기 위해, 독립변수가 1개인 단순회귀분석을 먼저 예로 든다. 독립변수인 단어의 빈도(X)만을 활용하여 종속변수인 텍스트의 난이도(Y)를 추정하는 경우를 생각해 보자. 10개의 텍스트를 표본으로 추출하여, X와 Y의 관계를 좌표에 나타낸 결과가 다음 <그림 4>와 같다고 가정하자.

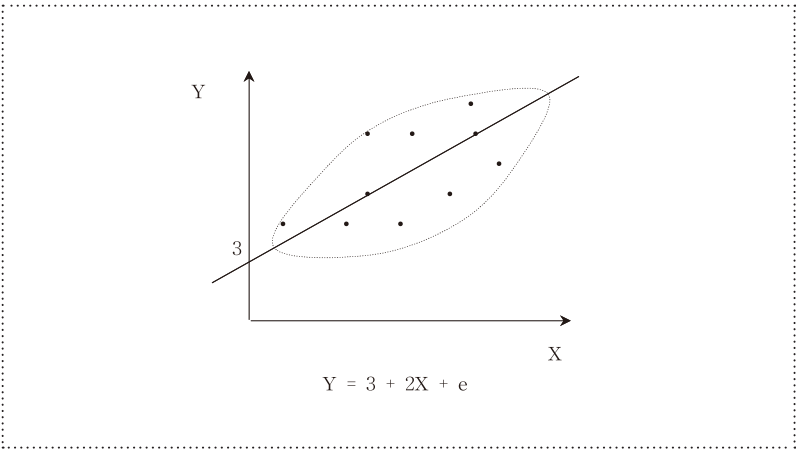


그림 4. 단순회귀분석의 선형함수 관계

위의 <그림 3>에서 각 점들은 단어의 빈도(X)에 대한 텍스트의 난이도(Y)를 나타낸다. 10개의 텍스트에 대하여 X와 Y의 관계를 가장 잘 나타내는 것은 타원의 가운데를 통과하는 직선이다. 회귀분석에서 선형함수 관계를 구하는 것은 바로 이 직선을 찾는 것을 말한다. e는 오차항이다. 오차항을 제외하고 위의 그림을 해석하면, 단어의 빈도(X)에 따른 의미적 복잡도가 0, 1, 2, ...로 올라감에 따라, 텍스트의 난이도(Y)는 3, 5, 7, ...로 증가한다. 즉, 텍스트의 난이도는 기울기인 2씩 증가한다. 따라서 이 직선을 찾게 되면, 새

로운 텍스트의 단어의 빈도(X)만으로 텍스트의 난이도(Y)를 예측할 수 있게 된다.

이제 단어의 빈도(X_1)와 문장의 길이(X_2)라는 2개의 독립변수를 가지고 텍스트의 난이도라는 종속변수(Y)를 예측하는 경우를 생각해 보자. 독립변수(X)가 2개이고 종속변수(Y)가 1개이므로, 이때에는 3차원의 공간에 나타난다. 위의 <그림 4>에서는 2차원의 공간에서 타원 모양이지만, 이제는 3차원의 공간상에서 럭비공 모양을 이루게 되는 것이다. 이 럭비공 모양의 한가운테를 자르면, 위의 그림처럼 타원 모양이 되고 마찬가지로 이 타원의 한가운테를 지나는 회귀방정식을 찾으면 된다. 이를 일반 회귀방정식으로 나타내면, $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + e$ 이다. 이 식은 텍스트의 단어의 빈도(X_1)와 문장의 길이(X_2)로 텍스트의 난이도(Y)를 예측한다. β_1 과 β_2 는 직선의 기울기이다. 기울기만큼 텍스트의 난이도가 높아진다.

Lexile의 경우에, 회귀분석을 통해 구한 회귀방정식은 다음과 같다 (Smith *et al.*, 1989: 10). 이 연구에서는 통계프로그램 SPSS를 활용하여, 문항의 지문의 난이도에 대한 회귀방정식을 찾는다.

$$Y = -3.23274 + 9.82247 \times \text{문장 길이} - 2.14634 \times \text{단어 빈도}$$

정리하면, 텍스트의 난이도를 정확하게 측정하기 위해서는 회귀분석을 활용해야 한다. 단어의 빈도와 문장의 길이를 독립변수로 하고, 텍스트의 난이도를 종속변수로 하는 회귀분석을 한다. 이러한 과정은 SPSS 통계 프로그램을 활용한다. 검증된 회귀방정식은 여러 가지 텍스트의 난이도를 분석하여, 독자에게 적합한 난이도의 텍스트를 제공하는 데 활용하게 된다.

3. 척도 변환

Lexile 연구자들은 자신의 연구 결과를 흔히 온도계에 비유한다. 온도계는 절대적인 척도로 되어 있는데, 척도가 절대적이기 때문에 언제, 어디에서든 활용할 수 있다는 것이다. 예를 들어, 아이의 체온이 높으면 병원에 데려가고, 바깥 기온이 너무 낮으면 외출을 삼간다. 마찬가지로, 독자의 읽기 능력과 텍스트의 난이도를 절대적인 척도로 나타낼 수 있다면, 그에 따른 교육적 의사 결정을 보다 쉽게 할 수 있다.

위에서 측정한 문항난이도와 텍스트의 난이도는 아직 상대적인 척도이다. 이를 절대적인 척도로 변환할 필요가 있다. 이 척도 변환이 세 번째 과제이다.

척도 변환을 위해서는 몇 가지 과정을 거쳐야 한다. 먼저, 척도의 기준점과 눈금을 정해야 한다. 절대적인 척도를 사용하는 온도계 역시 기준점과 눈금을 가지고 있다. 섭씨온도계의 경우, 기준점은 물의 어는점(0℃)과 끓는점(100℃)이며, 눈금은 1℃씩의 등간척도이다. Lexile의 경우, 가장 낮은 기준점은 기초 독본이고, 가장 높은 기준점은 전자백과사전이다. 이 기준점은 각각 1학년 2학기 교과서의 텍스트와 직장에서 성인이 접하는 텍스트에 해당한다. 그리고 나서 Lexile는 기초 독본과 전자백과사전 간의 차이를 1/1,000로 나타낸다. 즉, 두 기준점의 차이는 1,000L이다.

기준점을 설정하기 위해서, 초등학교 1학년 2학기 교과서와 대학 교재나 백과사전을 기준 텍스트로 활용하는 방향을 선택한다. 미국의 3학기제에서 1학년 2학기 교과서를 기준 텍스트로 정한 것은 1학년 1학기 교과서가 해독(decoding) 위주의 기초문식성 교재여서, 텍스트의 난이도를 분석하기가 어렵기 때문인 것으로 보인다. 우리나라도 1학년 1학기 교재는 기초문식성 지도 자료가 포함되어 있어 텍스트 난이도를 분석하기 어렵다.

Lexile에서 척도를 10단위로 정했다는 뜻은 10L씩 올라갈수록, 텍스트가 더 어렵다는 의미이다. 예를 들어, 900L의 텍스트보다 910L의 텍스트가

더 어렵다는 것이다. 그러나 김대희(2011: 162-163)에서 지적한 대로, 실제로 텍스트의 난이도를 그렇게 규정할 수 있는 것은 아니다. 10L 높은 텍스트가 실제로는 더 쉬울 수도 있다. 척도를 세분화할수록 정밀한 측정이 이루어지기도 하지만, 반대로 측정의 오차가 커질 가능성도 높아진다. 이를 기계에 비유하면, 어떤 기계의 부품들이 매우 작고 정교하다면 정밀한 작업을 수행할 수도 있지만, 반대로 고장이 나기도 쉽다. 따라서 척도의 간격을 10이 아니라 20~50 사이의 값으로 설정하여, 측정의 오차를 줄이는 방향으로 한다.

다음으로, 상대적인 척도를 절대적인 척도로 바꾸는 공식을 구해야 한다. 이 공식은 비교적 간단한 수학적 절차로 해결할 수 있다. 앞에서 기준점과 눈금을 정했다면, 상대적인 척도의 값을 절대적인 척도의 값으로 변환할 수 있다. Lexile의 경우에, 척도 변환의 공식은 다음과 같다(Smith *et al.*, 1989: 14).

$$(\text{Logit} + 3.3) \times 180 + 200 = \text{Lexile 난이도}$$

위의 식에서 Logit는 상대적인 척도의 난이도이고, 3.3과 180은 기준점 텍스트의 난이도 분석을 통해 얻은 상수이다. 이러한 공식을 구하여, 상대적인 척도를 절대적인 척도로 변환할 수 있다. 이와 같이, 기준점과 간격을 정하여, 척도 변환의 공식을 구한다.

정리하면, 세 번째 과제는 회귀방정식이 산출하는 상대적인 척도를 절대적인 척도로 변환하는 것이다. 이를 위해 먼저, 척도의 기준점과 눈금을 결정해야 한다. 기준점을 정하기 위해서, 초등학교 1학년 2학기 교과서와 백과사전을 기준 텍스트로 선정한다. 척도는 0~1,000을 범위로 하며, 눈금의 간격이 10 이상이 되도록 개발한다. 다음으로, 간단한 수학적 공식을 구하여, 상대적인 척도의 값을 절대적인 척도의 값으로 변환한다.

IV. 결론

이 논문의 목적은 독자의 읽기 능력과 텍스트의 난이도를 정확하게 측정하여, 독자에게 적합한 텍스트를 선정하는 방법의 개발 방향을 모색하는데 있다. 이를 위해 국내외의 주요 선행연구들을 분석한 후, 개발상의 주요 과제를 살펴보고 개발의 방향을 모색하였다.

국내의 연구들은 독자와 텍스트 측정에 관한 기존의 연구들을 비판적으로 고찰하거나, 이독성 공식에 기초하여 텍스트의 난이도를 측정하는 것들이었다. 2010년에 발표된 미국의 중핵성취기준은 텍스트의 복잡도를 성취기준의 하나로 제시하였다. 이 때, 독자에게 맞는 텍스트를 선정하는 양적 측정의 이론적 근거를 Lexile에 두고 있었다.

차세대 이독성 시스템이라 불리는 Lexile와 같은 방법을 개발하기 위해서는 문항 개발과 문항난이도 분석, 회귀분석, 척도 변환이라는 과제를 해결해야 한다. 선다형 독해 검사 문항을 개발한 후, 문항난이도를 분석하기 위하여 문항반응이론을 활용하고, 텍스트의 난이도를 측정하기 위하여 회귀분석 방법을 사용하며, 측정된 값을 절대적인 척도로 변환하는 것을 개발 방향으로 세웠다.

온도계가 날씨에 관한 모든 현상을 설명할 수 없는 것처럼, 텍스트 선정에 관한 양적 측정 방법도 읽기에 관한 모든 것을 설명할 수 없다(Smith *et al.*, 1989: 40). 그러나 이러한 측정 방법이 잘 개발된다면, 그 동안 교사나 전문가의 주관적 판단에만 의존해 왔던 독자에게 적합한 텍스트 선정의 문제를 해결하는 데 도움이 될 것이다.

연구의 초기 단계이므로 본문에서는 구체적인 개발 방향으로 언급하지 못했지만, 이 모든 과정이 자동화된 시스템으로 구축될 필요가 있다. Coh-Metrix 연구자들은 다음과 같이 진술했다. “20년 전에 우리의 동료들 중 학생의 글은 물론 전문가의 글을 등급화하고, 학생이 글을 잘 읽을 수 있도록

하는 컴퓨터 시스템에 내기를 거는 사람은 거의 없었다. 그러나 이러한 실제적인 목적을 달성하는 새로운 도구가 있다. 게다가, 이 시스템은 지금 공교육, 교과서 개발, e러닝에서 사용되는 지점에 이르고 있다.(Graesser *et al.*, 2011: 48).” 우리나라도 컴퓨터 언어학의 발달로 코퍼스가 구축되어 있고, 다양한 프로그램들이 개발되고 있으므로, 자동화 텍스트 분석 시스템의 구축이 가능할 것이라고 생각한다.

* 본 논문은 2013. 6. 26. 투고되었으며, 2013. 7. 8. 심사가 시작되어 2013. 7. 31. 심사가 종료되었음.

참고문헌

- 강상진(2003), 『회귀분석의 이해』, 서울: 교육과학사.
- 김광해(2003), 『등급별 국어교육용 어휘』, 서울: 박이정.
- 박순원(2011), 「독서지수의 활용 방안 연구」, 『새국어교육』 87, pp. 37-58, 한국국어교육학회.
- 성태제(2002), 『현대교육평가』, 서울: 학지사.
- 윤춘채 · 서혁(2010), 「표준화 독서 능력 검사도구 개발 연구 II」, 『독서연구』 24, pp. 289-312, 한국독서학회.
- 이성영(2011), 「초등 교과서의 이독성 비교 연구」, 『국어교육학연구』 41, pp. 169-192, 국어교육학회.
- 정혜승(2010), 「글 난이도(difficulty) 평가를 위한 질적 방법 연구」, 『국어교육』 131, pp. 523-549, 한국국어교육학회.
- 천경록(1995), 「기능, 전략, 능력의 개념 비교: 국어과 교육의 개념과 관련하여」, 『청람어문교육』 13(1), pp. 316-330, 청람어문학회.
- _____ (2006), 「독서 능력 표준화 검사 도구의 연구 개발」, 『독서연구』 15, pp. 407-436, 한국독서학회.
- _____ (2009), 「읽기 교육의 내용과 지식의 깊이」, 『독서연구』 21, pp. 319-348, 한국독서학회.
- 최숙기(2012), 「텍스트 복잡도 기반의 읽기 교육용 제재의 정합성 평가 모형 개발 연구」, 『국어교육』 139, pp. 451-490, 한국국어교육학회.
- 최영환(2003), 『국어교육학의 지향』, 서울: 삼지원.
- 한철우 · 이경화 · 최구홍(2007), 「유 · 초등 독서 능력 표준화 검사 도구 개발 연구」, 『독서연구』 18, pp. 321-354, 한국독서학회.
- Afflerbach, P., Pearson, P. D. & Paris, S. G.(2008), Clarifying difference between reading skills and reading strategies, *The Reading Teachers*, 61(5), pp. 364-373.
- Anderson, T. H. & Armbruster, B. B.(1986), Readable textbooks, or, selecting a textbooks is not like buying a pair of shoes. In Orasanu, J.(Ed). *Reading comprhension: from search to practice*, NJ: Lawrence Erlbaum Associates. Publishers.
- CCSSI(2010), *Common Core State Standards for English Language Arts & Literacy in History/Social Studies, Science, and Technical Subjects*, Washington, DC: CCSSO & NGA.
- Fry, E.(1968), A readability formula that saves time, *Journal of Reading* 11(7), pp. 265-271.
- Graesser, A. C., McNamara, D. S. & Louwerse, M. M(2011), Method of automated text analysis. In M. L. Kamil, P. D. Pearson, E. B. Moje & P. P. Afflerbach(Eds.), *Handbook of reading research*(Vol. IV), pp. 34-53. New York: Routledge.
- Hiebert, E. H.(2012), The Common Core State Standards and text complexity. *Teacher Librarian* 39(5), pp. 13-19.
- _____ (2013), Supporting students' movement up the staircase of text

- complexity, *The reading teachers* 66(6), pp. 459-468.
- Schnick T. & Knickelbine, M.(2000), *The lexile framework: An introduction for educators*, Durham, North Carolina: MetaMetrics, Inc.
- Smith, D. R., Stenner, A. J., Horabin, I., Smith, M.(1989), The Lexile scale in theory and practice: Final report for NIH Grant HD-19448. Durham, North Carolina: MetaMetrics, Inc.
- Spache, G. D.(1953), A new readability formula for primary-grade reading materials, *Elementary School Journal* 53, pp. 410-413.
- Walpole, S., Hayes, L. & Robnolt, V.(2006), Matching second graders to text: The utility of a group-administered comprehension measure, *Reading Research and Instruction* 46(1), pp. 1-22.
- <http://www.lexileframework.com>
- <http://cohmetrix.memphis.edu/>

독자에게 적합한 텍스트 선정 방법의 개발 방향

조용구

이 논문의 목적은 독자의 읽기 능력과 텍스트의 난이도를 정확하게 측정하여, 독자에게 적합한 텍스트를 선정하는 방법의 개발 방향을 모색하는데 있다.

독자에게 텍스트를 맞추기 위하여, 국내외의 주요 선행연구들을 분석한 후, Lexile와 같은 방법을 개발하는 데 초점을 두었다. 개발상의 주요 과제를 살펴보고 개발의 방향을 모색하였다.

차세대 이독성 시스템이라 불리는 Lexile와 같은 방법을 개발하기 위해서는 문항 개발과 문항난이도 분석, 회귀분석, 척도 변환이라는 과제를 해결해야 한다. 선다형으로 독해 검사 문항을 개발한 후, 문항난이도를 분석하기 위하여 문항반응이론을 활용하고, 텍스트의 난이도를 측정하기 위하여 회귀 분석방법을 사용하며, 절대적인 척도로 변환하는 것을 개발 방향으로 세웠다.

핵심어 독해, 문항반응이론, 이독성, 텍스트의 난이도, 회귀분석, Lexile

ABSTRACT

Developing Method of Matching reader to texts

Jo, Yong - gu

This thesis investigated the developing method of matching reader to texts. This thesis focused on the Lexile Framework which was developed by MetaMetrics.

After reviewing prior researches, this thesis presented three principles for developing of method of matching reader to texts. The method of matching reader to texts must be intended to 1) developing test of reading comprehension and analysing item difficulty, and 2) using regression analysis, and 3) translating the scale. To do this, this thesis suggested using item response theory, regression analysis and absolute scale.

KEYWORDS comprehension strategies, item response theory, readability, text complexity, regression analysis, Lexile