

인공지능 기반 한국어 시 생성 시스템 개발 연구

A Study on Korean Poetry Generation System Based on Artificial Intelligence

김 명 선 (Myung-sun Kim) CJ올리브네트웍스 AI연구소 연구원
정 우 혁 (Woo-Hyuk Jung) CJ올리브네트웍스 AI연구소 연구원
우 지 환 (Jihwan Woo) CJ올리브네트웍스 AI연구소장 / 고려대학교 기술경영대학원의 겸임교수, 교신저자

요 약

본 연구에서는 한국어 기반의 시를 창작하는 데 도움이 되는 문장들을 생성하는 인공지능을 개발하였다. 인공지능이 인간의 고유의 영역이라고 할 수 있는 창작에 대한 욕망과 창의력을 대신하는 것이 아니라, 인간이 창의력을 효율적으로 사용할 수 있도록, 창작에 밑바탕이 되는 문장들을 생성하는 데 초점을 맞추었다. 시인들과의 인터뷰를 통해서 8개의 다른 데이터 세트로부터 문장을 학습하여 8가지 장르의 시가 생성될 수 있도록 만들었다. 이 연구는 한국어를 활용한 문학 작품 생성 기술을 개발하였다는 점에서 차별성이 있으며, 이 연구를 확장해서 수필과 산문 또는 소설과 같은 다양한 형태의 문학 작품을 창작하는 데 도움이 될 수 있다는 점에서 큰 의미가 있다.

키워드 : 자연어생성, 시, 언어모델

I. 서 론

딥러닝을 기반으로 한 인공지능 기술의 발전으로 제조, 금융, 의료 등 다양한 분야에서 인공지능이 적용되고 있다(이승언 등, 2021; 천예은 등, 2021). 지금까지 여러 산업 분야에 적용되고 있는 인공지능 모델은 주어진 데이터를 학습하여 새로운 데이터가 주어졌을 때, 주어진 카테고리 중 하나로 분류를 하거나 또는 예측하는 역할을 주로 진행하였다. 예를 들어서, 동물이 있는 사진이 주어지면, 사진 안에 있는 동물이 개인지 고양이인지 분류를 하거나, 주식 가격 정보를 바탕으로 미래의 주식 가격을 예측하는 역할을 진행해 왔다.

그러나 GAN을 시작으로 새로운 데이터를 생성하는 인공지능 모델이 등장하면서, 인간의 창의성과 상상력을 극대화하는 문화 예술 분야까지 인공지능 모델의 사용이 확장되고 있다. 인공지능 기술을 활용해서 그림을 그리고, 글을 쓰고, 음악을 만드는 수준까지 도달한 것이다. 실제로 2022년도 8월에 열린 미국 콜로라도 주립 박람회 미술대회의 디지털 아트 부분에서는 ‘미드저니(Midjourney)’라는 인공지능 모델이 만든 ‘스페이스오페라극장’이 1등을 차지했다. 이러한 인공지능의 발전은 최근 OpenAI가 출시한 ‘chatGPT’(Susnjak, 2022)의 등장으로 본격적으로 주목을 받기 시작했다. 작사, 작곡, 시/소설 쓰기, 프로그램 코딩이나 리팩토링 같은

다양한 창작 활동을 높은 수준에서 시작할 수 있게 된 것이다. 인공지능을 창작의 수단으로 활용하는 방법들이 증가하고 있는 가운데, 이 논문에서는 인공지능을 활용하여 시를 생성하는 기술에 대한 연구를 진행하였다. 생성 모델을 개발하는 연구는 많이 진행되어 왔지만, 생성 모델을 활용하는 연구는 많이 부족한 현실이다. 문장을 창작하는 데 사용되는 많은 인공지능 모델이 영어로 되어 있는 가운데, 이 연구에서는 한글 데이터를 학습시켜서 시를 생성하였다는 점에서 한글을 이용한 문장 생성을 연구하는 데 도움이 될 것으로 판단된다. 또한 연구 결과는 현대 시인들과 협업을 통해서 출판까지 했다는 점에서 다른 연구와 차별점이 존재한다. 시의 경우 표현 형태와 길이 그리고 장르에 따라서 그 형태가 다양하다는 점이 존재한다. 시인들과 심도 깊은 인터뷰를 통해서 현대 시인들이 관심을 가지고 있는 8개의 장르(반전요소, 웹크롤링 데이터 기반, 순수 작품, 감정 풍부, 노래가사, 현대 시, 해외 소설)를 설정하고 데이터를 학습시켜서, 사용자의 취향에 따른 다양한 형태의 시가 생성될 수 있도록 하였다. 시에서는 짧은 길이에 의미를 함축해서 답아야 하기 때문에 단어의 선택적인 사용이 중요하다. 따라서 제외하는 단어 기능을 설정하여 이 단어가 시를 구성하는 문장을 생성하는 데 사용되지 않도록 제외하는 기능을 추가하였다. 이와 함께, 문학 작품 창작에 있어서 표현이 중요한 문제이기 때문에 학습에 사용된 문장들이 재 사용되어 표현이 될 수 있는 위험을 낮추기 위해서 학습 문장과 유사도를 계산하는 기능도 구현하였다.

이 연구에서 주목해야 할 점은, 시를 생성하는 기술에 대한 연구는 사람이 가진 창의력을 대신하는 것이 아니다. 창작에 사용되는 많은 시간이 다양한 아이디어들을 만들어 내는 과정에 있다는 점에 주목해서 인공지능은 다양한 문장들을 제시하고 시인들은 여기에 영감을 받아서 실제 창의력이 고도로 필요한 부분에 집중할 수 있도록 하는 데 있다는 것을 밝혀둔다. 본 연구는 시를 생성하는 인공지능 모델에 대해서 주목하였지만, 이 연구를

확장하면, 한국어를 기반으로 한 수필, 산문, 소설 등 다양한 형태의 문학 작품의 창작에 있어서 단서를 제공하는 모델을 만들 수 있다는 점에서 큰 의미가 있다.

II. 문헌연구

시는 상상력이 풍부한 문체로 생각과 감정을 표현하는 단어들의 연속이며, 일부는 엄격한 문학적 구문이나 형태를 따른다. 그것들은 예술적인 표현이고, 언어에 대한 깊은 지식과 숙달을 필요로 한다. 이와 같이 시 생성은 매우 어려운 작업으로 여겨지며 최근 10년 동안 많은 연구자들에게 관심을 끌고 있다. 시를 생성하기 위해 다양한 접근 방법이 시도되었는데, 본 연구에서는 이러한 방법론들을 크게 세 가지로 분류하여 검토하였다.

2.1 템플릿과 규칙 기반 방법론

운율, 강세, 단어 빈도와 같은 일련의 제약 조건에 따라 규칙을 구성하고 해당 규칙을 만족하는 텍스트를 생성하는 연구들이 진행되었다. Lamb and Brown(2019)에서는 시를 생성하기 위해 원본 텍스트에서 가능한 후보 절을 찾고 울감, 주제성, 형상화, 감정 측정에 대해 스코어링 해 점수를 높이는 방향으로 편집하는 시스템을 제시하였다.

2.2 최적화 기반 방법론

최적화 기법을 이용한 시 생성 연구 또한 활발히 진행되었다. Oliveira(2012)에 의하면 포르투갈의 시를 활용해 텍스트에서 총 700개의 패턴을 추출하고 진화 알고리즘을 통해 생성된 시에 대해 템플릿과의 음절 패턴을 비교한 점수를 이용하여 학습하는 방식으로 시를 생성하였다. Yan et al.(2013)에서는 중국 시를 생성하기 위한 생성 요약 프레임워크를 제시하였으며 시 구성 문제를 최적화 문제로 공식화하였다. 사용자가 지정한 의도를

전달하면 시스템은 말뭉치에서 후보 용어 검색 후 용어를 시의 형식에 맞게 음조와 리듬 요구사항을 충족하기 위해 최적화 프로세스를 수행한다. Manurung *et al.*(2012)에서는 의미성, 문법성, 시적 특성 3가지의 제약 조건을 충족하는 확률론적 시 생성 방법을 제안했다. 기존의 시를 통해 일반화된 템플릿을 정의하고 이를 충족시키는 목적함수를 최적화하도록 유전 알고리즘을 이용하였다.

2.3 딥러닝 기반 방법론

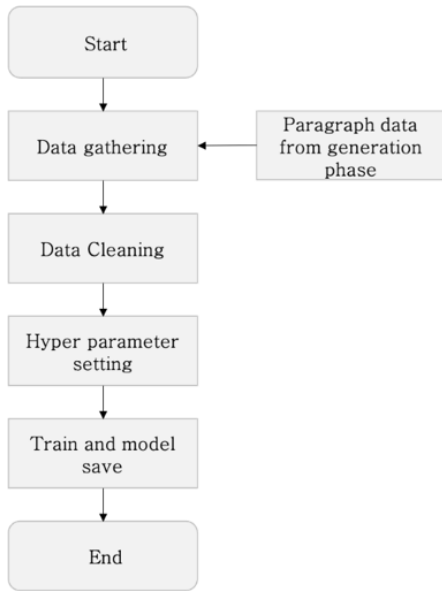
딥러닝 기반의 자연어 생성 분야의 비약적인 발전으로 시 생성 분야에 딥러닝을 활용하는 추세가 이어지고 있다. Zhang(2020)에 의하면 고대 시의 자동 생성 모델을 제안하였다. 원한 인코딩을 활용하여 시를 벡터로 변환하여 장단기 메모리(Long short-term memory: LSTM) 모델을 훈련시키고, 이중 언어 평가 언더스터디 점수(Bilingual Evaluation Understudy Score: BLEU)를 통한 자동 평가와 사람의 검증을 통한 수동 평가 2가지의 방법으로 생성된 시를 검증했다. Atassi(2022)에 의하면 아랍어 시를 생성하기 위해 LSTM, 양방향 장단기 메모리(Bi-Directional Long-Term Memory: Bi-LSTM) 및 게이트 순환 유닛(Gated Recurrent Unit: GRU) 활용하고 비교하였다. 입력 텍스트를 연속적으로 이동하여 소스와 타겟 쌍을 만들었으며 각각의 텍스트는 index로 벡터화되었다. 학습 결과 GRU 모델이 LSTM보다 성능이 우수한 것으로 나타났다. Hu and Sun(2020)에서는 중국 고전 시 생성을 위해 생성적 사전 훈련 변환기(Generative Pre-trained Transformer 2: GPT-2)를 활용하였다. 좀 더 통제된 생성 결과를 얻기 위해 고전 시를 문자 부분과 제목과 본문을 구분하는 구분자 부분으로 분리하여 각 부분에 가중치를 적용한 형태 변환 가중치 법(form-stressed weighting method)을 제안하였다. Beheitt and Hmida(2022)에서는 GPT-2를 활용한 아랍어 시 생성 방법을 제안하였다. 유창성, 일관성, 의미, 시성의 4가지 척도

로 사람이 생성 결과를 평가하였으며 BLEU 지표를 함께 확인하였다. Härmäläinen *et al.*(2022)에서는 프랑스어 현대 시 생성을 위해 강건히 최적화된 BERT 사전 교육 접근법(A Robustly Optimized BERT Pretraining Approach: RoBERTa)과 GPT-2를 활용하였다. 자연어 이해 성능과 생성 성능의 이점을 함께 취득하기 위해 모델의 인코더로는 RoBERTa가, 디코더로는 GPT-2가 사용되었다. 문법과 일관성 등의 6가지 평가 기준을 도입하여 사람이 매뉴얼하게 평가하였다. Santillan and Azcarraga(2020)에서는 트랜스포머를 활용하여 시를 생성하고 이를 문서 벡터 임베딩(doc2vec embedding)과 결합하여 자동으로 생성된 시를 평가하는 방법을 제안하였다. 생성된 결과는 임베딩 벡터의 코사인 유사도를 계산하여 유사한 정도를 확인하였다. 또한 특정 시인의 시를 개별적으로 학습한 결과를 normalized hit rate를 통해 평가하여 트랜스포머가 저명한 시의 스타일을 포착할 수 있고 생성 결과에 이를 반영하여 해당 시가 특정한 시인의 것으로 식별될 수 있다는 것을 확인하였다. Liu *et al.*(2018)에서는 3단계 멀티모달 중국어 시 생성 방법을 제안하였다. 주어진 그림에 대해 시의 첫 줄, 제목, 그리고 다른 줄들을 연속적으로 생성하는 데 계층-어텐션 시퀀스 투 시퀀스(hierarchy-attention seq2seq) 모델을 제안하였다. 생성 결과의 점수로 제목과 문단의 상관 계수를 Perplexity(PPL)와 곱한 형태로, 생성 결과의 합리성과 더불어 제목과의 유사성을 높이기 위한 점수를 고안하였다. 생성 결과의 척도는 BLEU와 RHYTHM 점수를 활용하여 생성된 시가 시로서의 운율과 규칙을 잘 준수하는지 평가하였다.

III. 제안 모델

3.1 학습 단계

학습 단계에서의 전체적인 flowchart는 <그림 1>과 같다.



〈그림 1〉 학습 단계 Flowchart

3.1.1 데이터 수집 및 전처리

본 연구에서는 서로 다른 시적 장르에 대하여 장르별 필요 데이터를 수집한다. 시의 장르는 총 3가지로 서정시, 소설, 산문시로 구분한다. 노래 가사와 서정시의 경우, 웹 크롤링을 통해 데이터를 수집하며, 소설과 산문시의 경우, 책의 내용을 광학 문자 인식(Optical Character Recognition, OCR) 기술을 통해 텍스트로 변환하여 수집한다. 시 장르별 수집 데이터 경로와 건수는 <표 1>과 같다.

〈표 1〉 장르별 학습 데이터 통계

장르	서정시	소설	산문시
샘플 수 (건)	4,208	979	1,362
샘플당 평균 글자 수 (자)	265	2,248	687

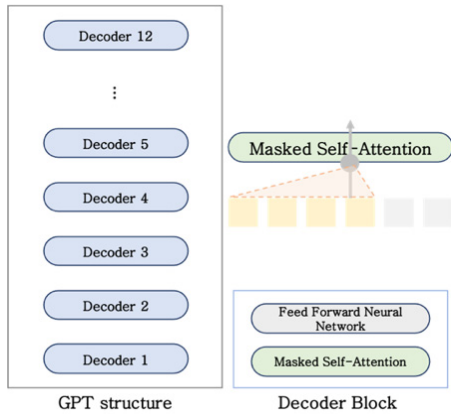
수집한 데이터셋에 대하여 정제를 진행한다. OCR을 통해 수집한 데이터의 경우, 변환 과정에서 문자의 오인식이 발생하며, 페이지 번호, 특수

문자 등 불필요한 문자가 포함 되어있다. 따라서, 이를 제거하기 위해 데이터 정제가 필수적이다. 또한 크롤링한 데이터의 경우 저자에 따라 글의 형식이 상이하여 이를 균일화하기 위한 절차가 필요하다.

순수 서정시와 산문시는 원본 데이터에서 제목을 분리한 뒤 불필요한 단어와 띄어쓰기를 제거한다. 그 후 행과 연의 구분을 위해 행 구분은 개행문자 1개, 연 구분은 개행문자 2개로 통일하여 규격화한다. 마지막으로 한 편의 시를 구분하기 위해 구분문자를 삽입한다. 구분문자는 ‘===’로 정의하며, 추론 시 해당 구분문자를 기준으로 추론 결과를 분리한다. 소설 장르는 서정시와 달리 제목이 포함되지 않아, 불필요한 단어와 띄어쓰기만 제거한다. 또한 대화가 주를 이루기 때문에 대화 구분을 위해 큰따옴표(“)를 삽입하고 개행문자를 삽입하여 자연스러운 대화문 형태로 통일한다. 마지막으로 한 편의 소설을 구분하기 위해 구분문자 ‘===’를 삽입한다.

3.1.2 모델링 단계

본 연구에서는 GPT-2 언어모델을 학습에 사용하며 전체적인 구조는 <그림 2>와 같다. GPT-3은 GPT-2에 비해 많은 데이터를 학습한 언어모델로써 GPT-2에 비해 다양한 어휘를 생성할 수 있는 장점이 있다. 그러나, 높은 성능의 컴퓨팅 파워가 요구되며, 학습 속도가 느린 단점이 있다. 비록 GPT-2가 GPT-3에 비해 적은 양의 데이터를 학습한 언어모델이지만, ‘시’와 같이 특정 하나의 분야에 대한 데이터를 활용하여 KoGPT-2를 파인튜닝 하는 경우, 맥락을 가진 문장을 생성하는 인공지능 모델을 구축할 수 있다. 또한, GPT-2는 GPT-3에 비해 파라미터 수가 적기 때문에 학습 속도가 빠르고 적은 컴퓨팅 파워가 요구된다 (Zong and Krishnamachari, 2022). 따라서, 본 연구에서는 pre-trained된 KoGPT-2 모델 기반으로 한국어 시 데이터를 이용하여 파인튜닝을 수행한다.



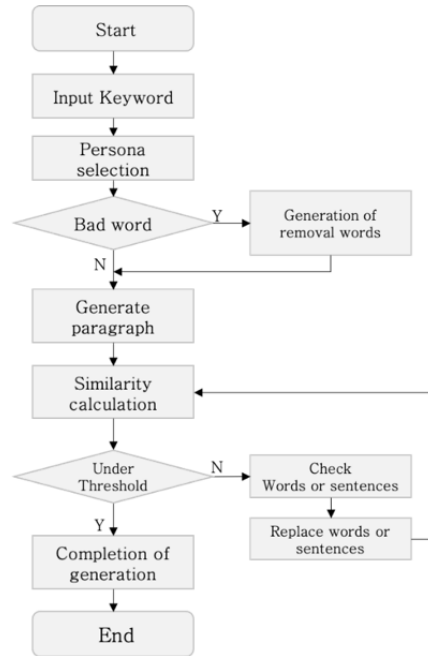
〈그림 2〉 GPT 모델 아키텍처

언어 모델은 문장의 각 단어에 확률을 할당하여 가장 확률이 높은 단어를 예측하여 반환한다. GPT 언어모델은 트랜스포머의 디코더 블록을 쌓아 올린 구조로 이루어져 있다. 입력 토큰들에 대해 다음에 올 토큰을 예측한 후 해당 토큰을 포함한 토큰들을 재입력하여 후행 토큰을 하나씩 생성하는 방식으로 동작 한다. 본 연구에서는 사전 학습모델로 SKT에서 배포한 한국어 GPT-2 모델을 이용한다. koGPT-2는 OpenAI에서 개발한 GPT 모델을 한국어 데이터셋으로 학습한 모델이다. 데이터셋에는 한국어 위키 백과, 뉴스, 모두의 말뭉치, 청와대 국민청원 등으로 구성되어 있다. 약 40GB의 한글 데이터셋으로 학습한 모델이며 GPT-2-small 모델로, 디코더 블록이 12개로 구성되어 있다. 딥러닝 모델의 하이퍼 파라미터 셋팅은 에폭(epoch)은 20, 기울기 감소율(weight decay)은 0.1, 학습률(learning rate)은 $1e-5$, 배치 사이즈(batch size)는 64이며, 학습 입력의 토큰(token) 개수는 256개로 설정한다. koGPT-2에 적용한 하이퍼 파라미터는 학습 데이터 수와 A100 GPU 리소스를 고려하여 실험적으로 설정하였다.

3.2 생성 단계

생성 단계에서의 전체적인 개략도는 <그림 3>

과 같다. 생성 단계에서는 주어진 키워드에 대해서 창작 문구를 생성되되, 원하지 않는 단어가 있는 경우 제외하는 단계를 거친다. 그 후 학습 데이터에 대해 과적합을 방지하기 위해 유사도 검사를 수행하며, 유사도가 0.8 미만인 경우에만 문장을 반환한다.



〈그림 3〉 생성 단계 Flowchart

3.2.1 제외 단어 설정 단계

본 연구에서 제안한 생성 모델은 사용자의 자유도를 높이기 위한 인공지능 모델이다. 사용자는 특정 단어 또는 문장을 활용하여 시를 생성할 수 있다. 뿐만 아니라 사용자가 원하지 않는 단어를 제외하고 시를 생성할 수 있다. 따라서 본 연구에서는 사용자의 의도를 반영할 수 있는 인공지능 모델로서 사용자가 원하지 않는 단어를 제외하고 싶은 경우, 제외 단어 모듈 규칙에 따라 해당 단어를 제외하고 완성된 글을 생성한다.

이때, 문장을 토큰나이징하는 방법에 따라 단어의 제외 여부가 달라지기 때문에, 가능한 모든

경우의 수를 반영하는 리스트 생성이 필요하다. 예를 들면 ‘자신’이라는 단어를 제외하기 위해 예상되는 제외 토큰들은 <표 2>와 같다. 단어가 임의의 조사와 조합되고, 이전 공백문자가 띄어쓰기나 개행문자냐에 따라 토큰나이징 결과가 달라지는 모습을 확인할 수 있다.

<표 2> 제외 토큰 시나리오 예시

문장	Tokenizing 결과	제외 token
내가 자신이 없어	_내가, _자신이, _없어	_자신이
자신 밖에 모르는	_자신, _밖에, _모르는	_자신
너는 왜\n자신밖에 몰라?	_너는, _왜, \n, 자신, _밖에, _몰라 _	자신
너는 왜\n자신이 없어?	너는, _왜, \n, 자, _신, _없어, ?	자
자신을 소중히 대하자	_자신을, _소중, 히, _대, 하자	_자신을

따라서 모든 경우의 수를 고려하기 위해 ‘은’, ‘는’, ‘이’, ‘가’, 등을 포함해 한국어 조사 427개로 구성된 리스트를 생성하고, 사용자로부터 입력받은 단어와 띄어쓰기, 그리고 조사 리스트를 조합한다.

예를 들면 ‘아침’이라는 단어를 제외하고자 할 때 가능한 token 조합은 _아침, _아, 아, 아침 4가지로 글자 수의 2배이다. 또한 여기서 조사가 붙는 경우 아침은, 아침이, _아침은, _아침이 등 조사 수의 2배만큼이 추가된다. 따라서 총 토큰 리스트의 수는 식 (1)과 같이 계산할 수 있다.

이 중에서 실제 단어정보에 등록되지 않은 토큰은 제외하고 남은 리스트를 반환하여 생성 시 제외한다.

$$F(x) = 2 \times N(x) + M \times 2 \quad (1)$$

x : 제외 단어, $N(x)$: x 의 글자 수

$F(x)$: x 에 대한 제외 토큰 후 보 수

M : 한국어 조사 수

3.2.2 생성 단계

주어진 키워드에 대해 문장 생성을 진행한다. 자기 회귀적 언어 생성 시, 단어 간 독립적이라는 가정하에 주어진 키워드에 대해 이후에 올 단어의 확률값은 식 (2)와 같이 조건부 확률의 곱집합으로 분해되어 계산한다.

$$P(w_{1:T} | W_0 = \sum_{t=1}^T P(w_1 | w_{1:t-1}, W_0), \quad (2)$$

$$w_{1:0} = \emptyset,$$

w_0 : 초기 키워드

디코딩 방법에 따라 생성 결과에 많은 차이가 있을 수 있기 때문에 강건한 디코딩 방법을 찾기 위해 하이퍼 파라미터를 조절한다. 생성 단계에서의 하이퍼 파라미터는 T 와 P 가 있다. 식 (2)에 의해 vocab 중 확률이 가장 높은 단어를 생성할 수 있으나, 여러 디코딩 방법론으로 생성 방법을 변형할 수 있다. Paulus et al.(2017)에서는 n-gram 페널티를 도입하여 n-gram이 2번 이상 반복되지 않도록 제약조건을 설정한다. 이는 시에서 같은 구절이 반복되는 것을 방지할 수 있다. 또한 단어의 확률 분포를 조절하는 방법으로 Temperature를 조절할 수 있다. Temperature는 식 (3)과 같이 소프트맥스(softmax) 기울기를 변화시킴으로써 확률 분포의 극단성을 조절한다. T 가 1에 가까울수록 확률 분포는 동일하게 유지되고 0에 가까울수록 기존에 확률이 큰 단어가 출력될 확률이 커진다.

$$P(x_i) = \text{softmax}(x_i / T) \quad (3)$$

T : temperature x_i : vocab of index i

샘플링을 통해 단어의 다양성을 제한할 수 있다. 일반적으로 많이 쓰이는 방법은 2가지로써, Top-K 샘플링과 Top-P 샘플링이 있다. Top-P 샘플링은 누적 확률 분포 값이 p 이내인 vocab 들만 선별하는 방법이다. 즉, 식 (4)를 만족하는 단어들만 선별한다.

$$\sum_{w \in V_{top-p}} P(w | W_0) = p \quad (4)$$

W_0 : context

p : threshold probability

Top-K 샘플링은 확률이 높은 k개의 단어를 샘플링해 단어의 분포를 제한하는 방법이다. 낮은 값으로 설정하면 선행 단어와 연관성이 낮은 단어가 나오는 것을 방지할 수 있다.

3.2.3 유사도 체크 단계

생성된 문장이 학습 데이터에 과적합 된 결과가 아니라는 것을 보장하기 위해 유사도 체크를 진행한다. 유사도 체크를 위해서는 문장들을 토큰나이저 후 벡터로 임베딩 하는 과정이 필요한데, 이 과정에서 문장 BERT(sentence BERT: sBERT) 모델을 사용한다(Reimers and Gurevych, 2019). sBERT는 각 문장에 대한 함축 정보를 통해 데이터를 학습하며, 각각의 문장 벡터 정보를 활용하여 문장 간 유사도를 분석하는데 효율적인 모델이다(김봉민, 박성배, 2020).

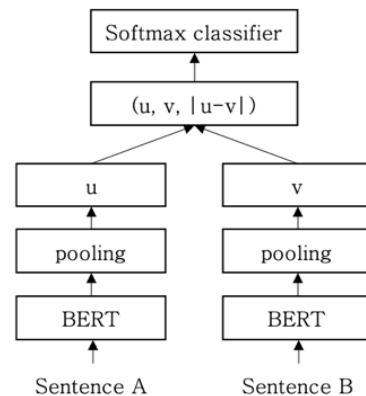
sBERT의 구조는 <그림 4>와 같다. u 벡터와 v 벡터를 각각 BERT의 입력으로 넣은 후 pooling을 통해 임베딩 벡터를 계산한다. 그 후 u 벡터와 v 벡터의 차이를 구한 후 세 가지 벡터를 연결한다. 마지막으로, 결합한 벡터에 소프트맥스를 적용하여 다중 클래스 분류 문제에 대한 학습을 진행한다. 카카오브레인의 KorNLU 데이터셋을 활용하여 학습된 사전 학습 모델을 활용하였으며, 추론 시 문장을 768차원의 벡터공간에 매핑한다. 그 후 임베딩 벡터에 대해 식 (5)를 통해 유사도를 계산한 뒤 유사도가 0.8 미만인 경우에만 문장을 반환한다.

본 연구에서는 유사도 체크를 위해 threshold 0.8을 적용한다. 생성 모델의 하이퍼 파라미터는 경험적 또는 실험적 반복을 통해 설정한다. 상기 언급된 바와 같이 생성 모델에서 temperature는 다음 단어 예측에 영향을 주는 대표적인 파라미터다. Temperature 값이 0인 경우, 다음에 생성할 후보

단어들의 확률값을 더 높일 수 있다. 이는 창의성을 낮추는 반면 문장의 맥락을 높일 수 있다. 반면에 Temperature 값이 1인 경우, 다음에 생성할 후보 단어들의 확률값을 비교적 균등하게 부여한다. 이는 창의적 글을 생성할 확률이 높아진 반면 문장의 맥락이 낮아질 수 있다. Temperature 값은 도메인별 경험에 따라 설정될 수 있으나, 생성한 문장의 맥락을 높이며, 창의성을 부여하기 위한 파라미터로는 0.8 또는 0.9를 적용한다. 이를 기반으로 유사성 체크에서도 문장의 맥락을 유지하고, 창의성을 높이기 위해 유사도 threshold를 0.8로 설정한다.

$$\cos(u, v) = \frac{u \cdot v}{\|u\| \|v\|} \quad (3.5)$$

u, v : 문장 A, 문장 B의 임베딩 벡터

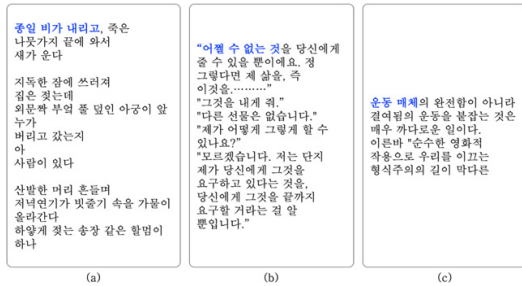


<그림 4> sBERT 모델 아키텍처

IV. 실험 및 결과

모델 결과 확인을 위해 3개의 서로 다른 키워드를 이용하여 문장을 생성한다. 키워드는 각각 서정시, 소설, 산문시의 검증 데이터셋에서 1개씩 선정한다. 사용된 키워드와 원본 데이터는 <그림 5>와 같다. 주어진 키워드에 대해 생성 결과를 <그림 6>~<그림 8>에 나타낸다. 전반적으로, 키워드의 소스와 관계없이 학습 데이터와 유사한 형태이다.

서정시 모델의 경우 개행문자로 행과 연이 구분되며 글의 생성 길이가 짧은 전형적인 서정시 스타일이 보이며, 소설의 경우 온점으로 구분되고 이야기의 기승전결이 뚜렷하며, 독백이나 이야기, 대화가 전개된다. 산문시의 경우에도 문장의 기승전결이 두드러지며 주로 서술하는 결과가 보인다.



〈그림 5〉 사용 키워드 및 원본데이터. (a) 서정시 (b) 소설 (c) 산문시

본 연구에서는 학습 데이터 수, 학습 시간 및 리소스 등의 문제로 최신 모델인 GPT-3, GPT-3.5 대신 GPT-2 기반의 모델을 활용한다. 제안 모델과 최신 모델인 GPT-3와 비교하기 위해 GPT-3 모델 학습을 진행한다. <그림 9>에서 확인할 수 있듯이, 생성 결과는 제안하는 GPT-2 기반 모델과 큰 차이가 없으며 동일한 키워드에 대해 비슷한 화풍의 결과가 생성되는 것을 확인할 수 있다.

본 연구에서 제안한 모델을 통해 생성한 결과가 원본의 문장과 얼마나 유사한지 확인하기 위해 원본 데이터와 유사도 점수를 계산한다. 유사도 계산 시 글자 수가 적으면 값이 크게 계산되기 때문에 가중평균을 계산하며, 그 값은 식 (6)을 통해 계산할 수 있다.

$$f(\vec{x}) = \sum_i \text{Score}_i \frac{\text{len}(x_i)}{\sum_j \text{len}(x_j)} \quad (6)$$

$\vec{x}$: 생성된 문장의 리스트

Score: sBERT 유사도

문장 단위 유사도 결과의 가중 평균값은 <표 3>과 같다. 편향을 제거하기 위해 딥러닝 모델의 시드(seed) 값을 변경하며 10번 생성한 결과를 평균 낸다. 서정시의 경우 대체로 높은 값을 확인할 수 있으며 소설은 대체로 낮은 값을 확인할 수 있다. 서정시는 상대적으로 정해진 형식과 단어의 반복이 이루어지나, 소설은 상대적으로 자유로운 형식과 다양한 주제로 전개되기 때문에 이러한 결과가 나타나는 것으로 생각된다.

〈표 3〉 유사도 점수의 가중평균

키워드	서정시	소설	산문시
종일 비가 내리고	80.5	77.7	70.5
어쩔 수 없는 것	72.4	70.5	71.9
운동 매체	70.6	66.8	63.2

본 연구에서는 결과 평가를 위해 PPL을 적용한다. Härmäläinen *et al.*(2022) 연구에서는 BLEU, Rhythm을 이용하여 모델 결과를 평가하였다. BLEU score는 일반적으로 번역에 많이 활용되는 평가 방법이다. Zhang(2020)의 연구에서는 생성한 테마별로 가장 많이 겹치는 단어가 있는 상위 20개의 첫 줄을 참고데이터로 설정하고, 생성한 첫 줄을 생성 데이터로 설정하여 BLEU 스코어를 계산하였다. 이러한 방법은 정답 셋이 아닌 생성한 데이터 간의 비교 score이기 때문에 객관적 평가가 어려울 수 있다. 또한, Rhythm은 성조에 대한 평가로 한국어에 적합하지 않다. 따라서, 문장의 헛갈리는 정도를 평가할 수 있는 PPL을 적용하여 생성한 문장의 결과를 평가한다. PPL은 생성 결과의 우도(likelihood)를 변형한 형태로, 앞선 맥락과 후행 단어를 비교했을 때 단어가 나타날 가능성이 얼마나 높은지 계산한 척도이다. 식 (7)을 통해 계산할 수 있다.

$$PPL(X) = \exp\left\{-\frac{1}{t} \sum_i \log p_\theta(x_i | x_{<i})\right\}, \quad (7)$$

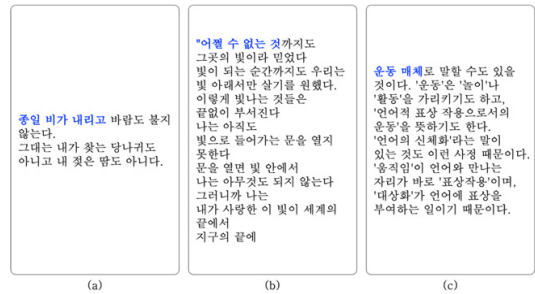
$\vec{x}$: 생성된 문장

생성 토큰을 256개 단위로 stride 하여 PPL의 평균값을 계산한 결과는 <표 4>와 같다.

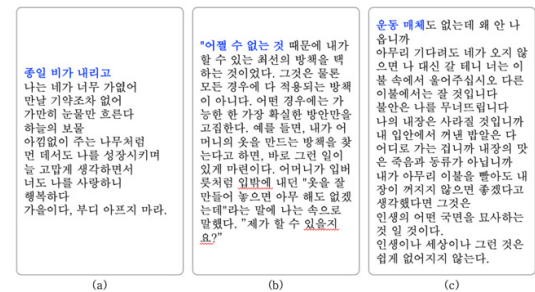
<표 4> 검증 데이터셋 PPL

	서정시	소설	산문시
PPL	57.8	73.6	56.5

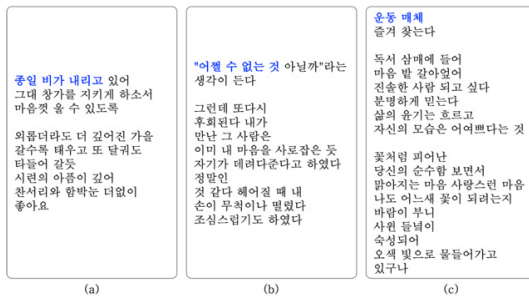
소설의 경우 PPL값이 높게 나타났으며 서정시와 산문시는 상대적으로 낮은 값이 확인된다. 소설은 서정시나 산문시와 달리 등장인물 간의 대화가 이어지고, <표 1>에서 확인할 수 있듯이 글의 길이가 길기 때문에 긴 문장이 생성될수록 맥락이 부자연스러운 경향이 있다. 또한 이야기의 맥락이 끝까지 유지되기 어려워 PPL 또한 높은 값이 얻어지는 것으로 생각된다. 유사도 점수와 PPL을 비교해 보았을 때 PPL이 높을수록 유사도 점수는 낮은 경향이 확인된다. 학습 데이터를 모방하지 않으면서 합리적인 생성 결과를 얻기 위해 두 점수를 비교하며 적절한 결과를 선택하는 방법이 중요할 것으로 여겨진다.



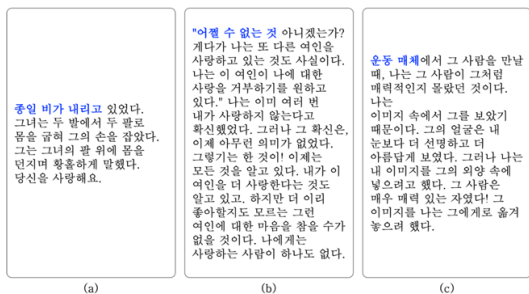
<그림 8> 산문시 모델 생성 결과



<그림 9> GPT-3 기반 생성 결과



<그림 6> 서정시 모델 생성 결과



<그림 7> 소설 모델 생성 결과

생성 결과에 대해 총 9명의 시인과 20명의 일반 사용자에게 설문조사를 진행한다. 시인 대상의 설문에서는 긍정적인 평가를 진행하였으며 일반 사용자를 대상으로는 정량적인 평가를 진행한다.

시인을 대상으로 한 평가에서 생성 결과의 장점으로는 시 다운 글이 나오지 않고, 평소에 사용하지 않는 단어/문장이 나와서 생각의 전환이 된다는 의견이 있다. 또한 처음 주제 설정, 작법에 따라 다른 결과물이 나온다는 의견이 있다. 단점으로는 기계적이고 예상되는 글이 생성되고, 익숙한 한국어 조합이라는 의견이 있다. 또한 많은 결과값이 '사랑'으로 귀결되어 피곤함을 느꼈다는 의견도 있다. 결론적으로 대부분의 시를 시인이 직접 써야 하나 새로운 키워드를 토대로 시발점이 된다는 의견이다.

<표 5>는 생성 결과의 장르별 사용량 및 선호도의 통계이다. 사용량은 산문시, 서정시, 소설 순으로 많았으며 선호도도 동일한 순서로 나타난다.

대체로 산문시를 가장 선호하고 소설에 대한 선호도는 낮은 것을 확인할 수 있다. 소설의 경우 긴 문장에 걸쳐 이야기가 진행되기 때문에 한 문단 정도의 짧은 생성 결과로는 많은 내용을 담을 수 없어 선호도가 낮은 것으로 해석된다. 또한 서정시에 비해 산문시가 높은 선호도를 보이는 이유는 시인은 새롭고 시다운 글이 나오지 않는 것을 선호하므로 나타나는 결과로 판단된다. 이 점은 시인과 일반 사용자의 가장 큰 차이점으로 보인다.

〈표 5〉 장르별 사용량 및 선호도

	사용량	선호도
서정시	35%	25%
소설	10%	0%
산문시	54%	75%

일반 사용자를 대상으로 한 설문은 총 5가지를 준비한다. 설문 내용은 기존 관련 연구에서 수행되었던 설문지를 참조하였다(Hämäläinen *et al.*, 2022). (1) 인공지능 생성 시는 사람이 만든 시에 비교하면 일반적이다. (2) 인공지능 생성 시는 이해하기 쉽다. (3) 인공지능 생성 시는 문법적으로 정확하다. (4) 인공지능 생성 시는 영감을 준다. (5) 인공지능 생성 시는 감수성을 불러일으킨다.

〈표 6〉 일반인 대상 AI 생성 결과 만족도 평가

	Q1	Q2	Q3	Q4	Q5
평균	3.28	3.56	3.31	3.76	3.50
표준편차	0.61	0.76	0.46	0.64	0.83

〈표 6〉은 설문에 대한 평균과 표준편차를 나타낸다. 가장 높은 점수를 받은 설문은 인공지능 생성 시가 영감을 준다는 결과로, 시인들을 대상으로 한 인터뷰 결과와 일관성 있게 나타났다. 또한 가장 낮은 점수를 받은 설문은 ‘인공지능 생성 시가 일반적이다’라는 설문이다. 시인들 또한 인공지능 생성 시가 기계적이라는 의견이 있으므로, 일반적이지 생성 결과 말투가 사람과 다른 점이

점수가 낮게 나온 원인으로 보인다.

일반 사용자와 시인을 대상으로 한 평가를 바탕으로 제안 모델의 개선점을 정리하면 다음과 같다. 제한된 데이터와 장르로 인해 많은 결과값이 비슷한 표현이 중복되는 점을 고려할 때, 다양한 장르의 작품을 추가하는 개선 방향이 있다. 또한 시 이외에 심리학, 미술 분야, 자연/식품 도감과 같은 예상치 못한 데이터를 추가하는 방법이 있다.

V. 결 론

5.1 요약

본 연구에서는 인간의 고유 영역인 예술의 창작 영역에 AI가 어떻게 접목될 수 있는지 알아보기 위해 AI 언어모델 기반 시 생성 시스템을 제안하였다. 작가들과 심층 인터뷰를 통해서 서정시, 산문시, 현대 시에 대해 데이터 수집 및 정제, 학습을 진행하였고 학습 결과에 대해 sBERT를 통한 원본 데이터와의 유사도 검증을 진행하여 표절을 방지하였다. 학습 데이터의 문학적 특성에 따라서, 사용되는 단어와 표현 방법이 상이하기 때문에 모델을 학습하기 어려운 문제가 있었으나, 이를 측정할 수 있는 지표(Measure)인 PPL을 제안하였다. 학습 및 생성 결과 소설은 유사도 값이 낮은 반면 PPL이 높은 경향을 확인했고, 서정시는 유사도 값이 높은 반면 PPL이 낮은 경향을 확인했다. 이는 소설 모델의 생성 결과가 맥락이 유지되는 결과를 생성하지 못하는 반면 원본 데이터와의 중복성은 낮고, 서정시 모델의 생성 결과는 맥락이 유지되는 반면 원본 데이터와 중복성이 높다고 판단된다. 따라서 두 지표를 적절히 판단한 결과를 생성하는 모델에 대해 취사선택이 필요하다. 또한 인공지능과 작가와의 협력에 대한 모델을 최초로 제시하고 창작 결과에 대해 분석을 진행하였다. 창작물의 성능을 개선하기 위해서 새롭게 학습되어야 할 데이터를 분석하여 제시하였다. 본 연구에서 제안한 방법은 시뿐만 아니라 웹 소설, 카피 라이

트 등 여러 산업에 활용 가능할 것으로 판단되므로 다양한 도메인에 적용이 기대되므로 상업적인 가치가 있다고 판단된다.

5.2 학술적 함의

본 연구가 제시하는 학술적인 시사점은 다음과 같다. 본 연구는 생성형 AI 기술을 작가들이 활용하여 창의성을 높일 수 있는 방법을 제시하였다. 기존의 연구는 컴퓨터 공학 분야에서 언어 모델을 정교하게 만드는 것에 집중되어 있었다(Beheitt and Hmida, 2022; Hu and Sun, 2020). 본 연구는 다양한 문장 생성 방법 중, 시가 가지는 특수성인 운율, 글자 수 제약, 단어의 함축 및 한국어의 특수성을 고려하여 글을 생성하는 알고리즘을 제안하였다는 점에서 기존 연구와 차별성을 보이며 학술적인 함의를 가진다. 또한, 본 연구는 생성된 문장들이 학습 데이터에 과적합 되지 않았음을 보장하기 위한 유사도 체크 방법을 도입하였습니다. 따라서 유사도 체크와 하이퍼 파라미터 설정을 통해 생성된 문장의 다양성과 창의성을 유지하면서도 학습 데이터에 과적합 되지 않도록 한 점에서 학술적인 함의를 내포하고 있다. 학술적인 함의 측면에서 본 연구는 인공지능을 활용한 문학 창작에 대한 이해를 높이고, 창의적인 아이디어 생성 과정에서 인공지능의 역할과 한계를 명확히 밝혀냈다. 인공지능은 창작 과정에 참여하는 도구로서의 역할을 수행하며, 시인들이 현실적으로 어려움을 겪을 수 있는 시간과 노력 소요가 큰 부분에 대한 지원을 제공할 수 있다. 이를 통해 사람과 인공지능의 협업이 창의적인 문학 작품 창작을 더욱 풍부하고 다양한 방향으로 이끌어낼 수 있음을 이론적으로 입증하였다.

5.3 실무적 함의

초거대 언어모델(LLM; Large Language Model)을 발전에 따라 산업 전반에서 AI를 활용한 글을

창작하는 일들은 활발히 진행될 것으로 판단된다. 본 연구는 창작자들이 AI를 활용해서 글을 창작하기 위해 필요한 데이터 수집, 전처리, 모델링 및 창작한 결과물의 표절 방지를 위한 전체 프로세스를 제시하였다. 초거대 모델이 파라미터 수를 늘려가면서 경쟁하는 가운데, 파라미터 수가 적은 모델을 활용해서도 충분히 좋은 결과를 제시할 수 있음을 보여 주었다. 실무에서는 파라미터 수가 증가하면 학습에 더 많은 시간과 비용이 필요하기 때문에, 적은 파라미터를 유지하는 것이 중요하다. 제안된 연구를 활용해서, ‘콘텐츠 생성 및 마케팅’, ‘개인화된 커뮤니케이션’, ‘문서 작성 및 요약’, ‘자동 글쓰기 도구’ 등을 제작할 수 있다. AI 모델을 활용하여 시를 작성하는 방법은 시인이나 문학 작가에게 창의적인 영감을 줄 수 있다. AI 모델은 다양한 문체와 주제를 다룰 수 있으며, 예상치 못한 시적 표현이나 아이디어를 제공할 수 있다. 이를 통해 시 작성에 대한 창의성과 다양성을 높일 수 있게 된다. AI 모델을 활용하여 문학적인 작품을 생성하고 편집하는 데에도 활용할 수 있다. 소설, 수필, 단편 등 다양한 형식의 문학 작품을 자동으로 생성하거나, 기존 작품을 수정하고 발전시킬 수 있는 것이다. 이를 통해 문학 작품의 창작과 편집 과정에서 새로운 아이디어와 문체를 탐색할 수 있는 장점이 존재한다.

기술적인 측면에서도 실무적 함의들이 다수 존재한다. 첫 번째는 과적합의 방지 기능이다. 유사도 체크를 통해 생성된 문장이 학습 데이터에 과적합 되지 않도록 시스템을 구성하였다. 학습 데이터에 대한 유사도 분석을 통해 이미 존재하는 문장과 유사도를 계산하고, 설정한 임계값에 따라 유사도가 낮은 문장을 선택하거나 생성할 수 있다. 이를 통해 생성된 문장의 다양성과 새로운 아이디어의 도출을 장려할 수 있게 된다. 두 번째는 창의성을 유지하는 측면이다. 하이퍼 파라미터 설정에서 ‘Temperature’ 값을 조절하여 창의성을 유지할 수 있게 만들었다는 점이다. 값이 낮을수록 예측된 단어의 확률이 높아져 보수적이고 일관

된 문장을 생성하게 되고, 값이 높을수록 예측된 단어의 확률이 균등하게 부여되어 창의적이고 다양한 문장을 생성할 수 있게 된다. 이를 통해서 문장 생성의 다양성과 창의성을 조절할 수 있다. 마지막으로는 자동 문자 평가 기능이다. 유사도 체크와 함께 다양한 자동 문장 평가 지표를 활용하여 생성된 문장의 질을 평가할 수 있다. 생성된 문장과 실제 문장 사이의 유사성을 측정하거나, 언어 모델의 편향성을 평가하는 지표를 도입하여 문장의 품질을 평가하고 개선할 수 있기 때문에 다양한 방향으로 활용 가능하다.

5.4 향후 연구 제언

본 연구에서 부족한 점들은 다음 연구 주제로 제언하며 연구가 발전되기를 기대한다. 우선, 다양한 작가 스타일의 모델 개발이 필요하다. 현재 제안된 시 생성 시스템은 다양한 작품을 생성할 수 있지만, 작가의 개별적인 스타일을 더 잘 반영할 수 있는 모델 개발이 필요하다. 예를 들어, 특정 작가의 작품을 기반으로 한 스타일 트랜스퍼 모델을 개발하여 해당 작가의 특징을 더욱 정교하게 재현할 수 있다. 다음으로는 대화형 창작 시스템의 연구를 제안한다. 단순히 작품을 생성하는 것뿐만 아니라, 작가와 독자 혹은 사용자 간의 상호작용을 통해 대화형 창작 시스템을 구축하는 것도 본 연구를 발전시킬 수 있을 것으로 기대된다. 작가와 독자 간의 상호작용을 통해 작품의 흐름이 변화하거나 독자의 요구에 따라 작품이 개인화되는 기능을 개발하는 것이 가능하기 때문이다. 마지막으로 작품 평가 및 피드백 시스템에 관한 연구이다. 작품의 품질을 평가하고 피드백을 제공하는 시스템을 개발하는 것도 중요한 연구 주제이다. 이러한 연구를 통해서 효과적인 평가 척도를 개발하고, 사용자의 의견을 수집하고 분석하여 작품의 개선을 도모하는 방법을 탐구할 수 있을 것으로 기대한다.

사실 향후 연구로 제안한 위의 주제들은 시를

중심으로 한 예술 창작에 초점을 맞춘 것이지만, 다른 장르나 도메인에도 적용 가능하다. 이러한 연구를 통해 인공지능과 예술의 접목이 더욱 깊어지고 다양한 형태로 발전할 수 있을 것으로 기대된다.

본 연구는 인공지능을 활용한 문학 창작에 대한 새로운 지평을 열었으며, 다양한 연구와 협업을 통해 미래의 문학 창작에 새로운 가능성을 제시하였다. 이를 바탕으로, 우리는 문학의 영역에서 인간과 인공지능의 상호작용을 통해 창의적이고 진보한 작품을 창조해 나갈 수 있을 것이다.

참고 문헌

- [1] 김봉민, 박성배, “한국어 문장 유사도 측정을 위한 Sentence BERT”, *한국정보과학회학술대회*, 제47권 제2호, 2020, pp. 1376-1378.
- [2] 이승언, 김현세, 김남호, 엄정윤, 우지환, “딥러닝 기술을 활용한 금융거래 실소유주 구별에 대한 연구”, *한국데이터정보과학회지*, 제32권, 제4호, 2021, pp. 781-797.
- [3] 천예은, 김세빈, 이자윤, 우지환, “설명 가능한 AI 기술을 활용한 신용평가 모형에 대한 연구”, *한국데이터정보과학회지*, 제32권, 제2호, 2021, pp. 283-295.
- [4] Atassi, A. and I. El Azami, “Comparison and generation of a poem in arabic language using the lstm, bilstm and gru”, *Journal of Management Information & Decision Sciences*, 25, 2022.
- [5] Beheitt, M. E. G. and M. B. H. Hmida, “Automatic arabic poem generation with gpt-2”, *In ICAART*, Vol.2, 2022, pp. 366-374.
- [6] Oliveira, H. G., “Poetryme: A versatile platform for poetry generation”, *Computational Creativity, Concept Invention, and General Intelligence*, Vol.1, 2012, Article 21.
- [7] Hämmäläinen, M., K. Alnajjar, and T. Poibeau, “Modern french poetry generation with roberta

- and gpt-2”, arXiv preprint arXiv:2212.02911, 2022.
- [8] Hu, J. and M. Sun, “Generating major types of chinese classical poetry in a uniformed framework”, arXiv preprint arXiv:2003.11528, 2020.
- [9] Lamb, C. and D. G. Brown, “Twitsong 3.0: Towards semantic revisions in computational poetry”, In *ICCC*, 2019, pp. 212-219.
- [10] Liu, D., Q. Guo, W. Li, and J. Lv, “A multi-modal chinese poetry generation model”, In *2018 International Joint Conference on Neural Networks (IJCNN)*, 2018, pp. 1-8.
- [11] Manurung, R., G. Ritchie, and H. Thompson, “Using genetic algorithms to create meaningful poetic text”, *Journal of Experimental & Theoretical Artificial Intelligence*, Vol.24, No.1, 2012, pp. 43-64.
- [12] Paulus, R., C. Xiong, and R. Socher, “A deep reinforced model for abstractive summarization”, arXiv preprint arXiv:1705.04304, 2017.
- [13] Reimers, N. and I. Gurevych, “Sentence-bert: Sentence embeddings using siamese bert- networks”, arXiv preprint arXiv:1908.10084, 2019.
- [14] Santillan, M. C. and A. P. Azcarraga, “Poem generation using transformers and doc2vec embeddings”, In *2020 International Joint Conference on Neural Networks (IJCNN)*, 2020, pp. 1-7.
- [15] Susnjak, T., *Chatgpt: The end of online exam integrity?*, 2022.
- [16] Yan, R., H. Jiang, M. Lapata, S.-D. Lin, X. Lv, and X. Li, “Poet: Automatic chinese poetry composition through a generative summarization framework under constrained optimization”, In *Twenty-Third International Joint Conference on Artificial Intelligence*, 2013.
- [17] Zhang, H. and Z. Zhang, “Automatic generation method of ancient poetry based on lstm”, In *2020 15th IEEE Conference on Industrial Electronics and Applications (ICIEA)*, 2020, pp. 95-99.
- [18] Zong, M. and B. Krishnamachari, “A survey on GPT-3”, arXiv preprint arXiv:2212.00857, 2022.

A Study on Korean Poetry Generation System Based on Artificial Intelligence

Myung-sun Kim^{*} · Woo-Hyuk Jung^{*} · Jihwan Woo^{**}

Abstract

In this study, we developed an AI-based system to generate sentences that assist in creating Korean poetry. Instead of replacing the creative aspect of composition, which is considered a unique domain of humans, the focus was on generating foundational sentences to enhance human imagination efficiently. By conducting interviews with poets, the researchers extracted sentences from eight distinct datasets, enabling the generation of poetry across eight different genres. This study stands out for its innovation in developing a method for crafting literary works in Korean. Its significance lies in its potential to facilitate the creation of diverse literary forms such as essays, prose, or novels.

Keywords: *Poetry Generation Model, Artificial Intelligence-Based Language Model, and Korean Sentence Generation Model*

^{*} Researcher, CJ OliveNetworks AI Research

^{**} Corresponding Author, CJ OliveNetworks AI Research, Korea University School of Management of Technology

○ 저 자 소 개 ○



김 명 선 (ms.kim83@cj.net)

광주과학기술원에서 에너지 데이터 분석을 전공하고, 전기전자컴퓨터공학 석사 학위를 취득하였다. 지능 정보 시스템 연구실에서 데이터 분석 관련 연구들을 진행하였으며, 현재 CJ올리브네트웍스 AI연구소에서 재직 중이다. 주요 연구 분야는 시계열 데이터 분석이다.



정 우 혁 (wh.jung2@cj.net)

현재 CJ올리브네트웍스 AI연구소에서 Natural language processing 및 데이터 분석 연구를 수행중이다. 주요 연구 분야로는 NLP, Data Analytics, HCI, Bio signal processing 등이다.



우 지 환 (jihwan_woo@korea.ac.kr)

삼성전자 리서치 책임연구원, Carnegie Mellon University Robotics Institute 방문연구원으로 근무한 후, 현재 CJ올리브네트웍스 AI연구소장 및 고려대학교 기술경영대학원의 겸임교수로 재직중이다. 주요 연구분야는 AI, Computer Vision 및 AI기반의 디지털 트랜스포메이션의 산업 적용이다. 관련하여 20여 건의 국내외 특허 등록 및 10여 편의 논문을 게재하였다.

논문접수일 : 2023년 03월 07일

게재확정일 : 2023년 06월 12일

1차 수정일 : 2023년 06월 01일