

자료집

2020 한국지능정보시스템학회
춘 계 학 술 대 회

포스트 팬데믹 시대의 디지털 비즈니스

2020년 6월 12일(금)

온라인 및 현장 (연세대학교 백양로 B1 김순전홀)

주 최: (사)한국지능정보시스템학회

후 원:  i-Scream edu  KTNET  LSDTech  MarkAny*  아이티센
한국과학기술정보연구원

 EUGENE
유진테크서비스

 GKeS (주)지케스

 KSTEC  n3n
클라우드



한국지능정보시스템학회
Korea Intelligent Information Systems Society

디지털 무역·물류 혁신 플랫폼 기업, 한국무역정보통신

- (주)한국무역정보통신(KTNET)은 한국무역협회의 100% 출자로 설립된 ICT 중견기업으로서 무역·물류업계의 경쟁력을 강화하는데 중추적 역할을 하고 있습니다.
- 앞으로도 블록체인과 빅데이터, AI 등 4차 산업혁명 기술을 접목하여 대한민국의 디지털 무역·물류 발전에 기여하겠습니다.



LSDTech



MAC-T

Multi Array Channel -Transaction

서버 병목현상 제로화 기술을 적용한
고성능의 서버를 경험하십시오

- ▶ 전기요금 40% 절감!!
- ▶ 딥러닝/머신러닝/렌더링 처리시간 단축!!

엘에스디테크(주) 1599-7230

서울 금천구 벚꽃로 234, 1904호(가산동,에이스하이엔드타워6차)
lsdtech@lsdtech.co.kr www.monstertech.co.kr

기술에 기술을 더하는⁺ 새로운 마크애니



MarkAny*

대표전화 02-2262-5222
04615 서울시 중구 퇴계로 286 쌍림빌딩 10층

당신이 오늘 만날 수 있는 혁신

Innovation & Reality

4차산업 플랫폼
비즈니스 전문그룹

cen 아이티센그룹
ITCENGROU

Strategic Digitization

High-technology and High Quality IT Service

유진 IT 서비스는 기술과 품질, 고객의 부가가치를 지향합니다



EuTram

지능형 유통 솔루션

중소 유통업체의 체계적인 상품정보관리, 손쉬운 판매 채널 확대, 효율적인 재고관리를 제공하는 지능형 유통 솔루션

EuPMS

프로젝트 관리 솔루션

유진IT서비스가 보유한 EuSM의 프로젝트 관리방법론을 기반으로 구현된 프로젝트 관리도구

NBP Cloud

클라우드 플랫폼 서비스

네이버, LINE, BAND 등 다양한 서비스 운영 경험을 통해 검증된 클라우드 서비스

Postgre SQL EDB

오픈소스 DBMS

가장 안정적인 오픈소스 DBMS
평가받는 엔터프라이즈급 데이터 베이스 시스템

Dell EMC

플랫폼 장비 공급

Dell EMC가 제공하는 혁신적인 서버와 스토리지로 IT혁신을 가속화



공공와이파이 통합 관리 솔루션

와이파이지킴-e

와이파이지킴-e는 AP의 설치 위치 및 오류정보를 한눈에 볼 수 있는 대시보드를 제공하며 수많은 AP기기를 통합해서 관리할 수 있도록 도와주는 서비스입니다

와이파이지킴-e의 레퍼런스

- 1 공공와이파이 통합관리 시스템 구축 (2018-2020)
- 2 학교와이파이 관리 시스템 1차~4차 구축 (2019-2020)
- 3 지방자치단체 와이파이 관리시스템 구축 (2018-2020)



인공지능



축적된 경험



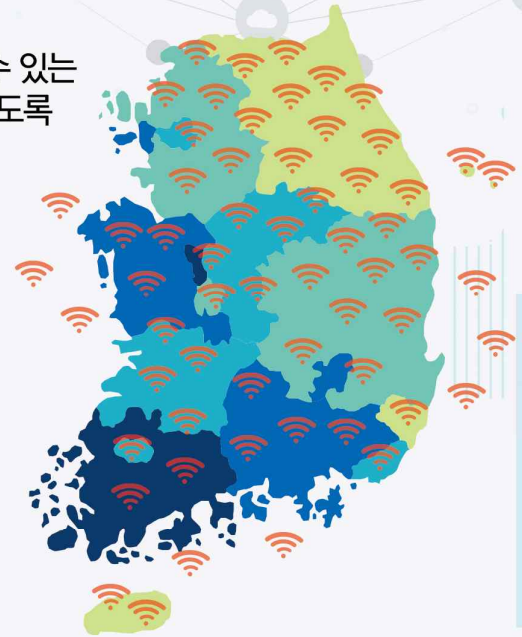
인력 절감



이용자 편의



W/H벤더 연계



총선 공약 1호 공공와이파이 설치 사업

HOT ISSUE

정부는 2022년까지 전국에 5만 3천여개의 공공와이파이를 확대 구축할 계획이라고 밝혔습니다. 공공와이파이 구축 공약에 드는 비용은 3년간 총 5780억원으로 추정하고 있습니다.



지능형 모니터링
클라우드 서비스

IT 운영 전담 인력이 부족한 기업들의 IT인프라를 지킴-e 클라우드 모니터링 센터에서 365일 24시간 원격으로 관제하는 서비스입니다.

**2020 클라우드서비스 지원제도로
정부 지원금 70% 받아 운영하세요**



저렴한 비용



맞춤형 관제



운영 및 교육

지케스와 함께할 인재를 찾습니다

신입/경력사원 모집

연구개발	0명	업무: wNMS, SMS, EMS, 관제시스템 연구개발 경력: 2년~ 10년(신입가능)
통합관리센터 운영요원	0명	업무: 통합관리센터 운영 민원대응 자격: 초급 자격요건 이상
솔루션 구축 유지 보수	0명	업무: 당사 솔루션 구축 및 유지 보수 경력: IT관련 업무자
영업사업부	0명	업무: 사업기획, 사업발굴, 사업관리 경력: 2년~ 10년(신입가능)

지원방법: 이메일 접수 recruit@gkes.co.kr



RPA 솔루션

Robotic Process Automation

소프트웨어 로봇을 통한 업무 프로세스 자동화
반복적/규칙적인 대량 업무 프로세스 자동화
편리한 GUI를 통한 자동화 BOT 생성



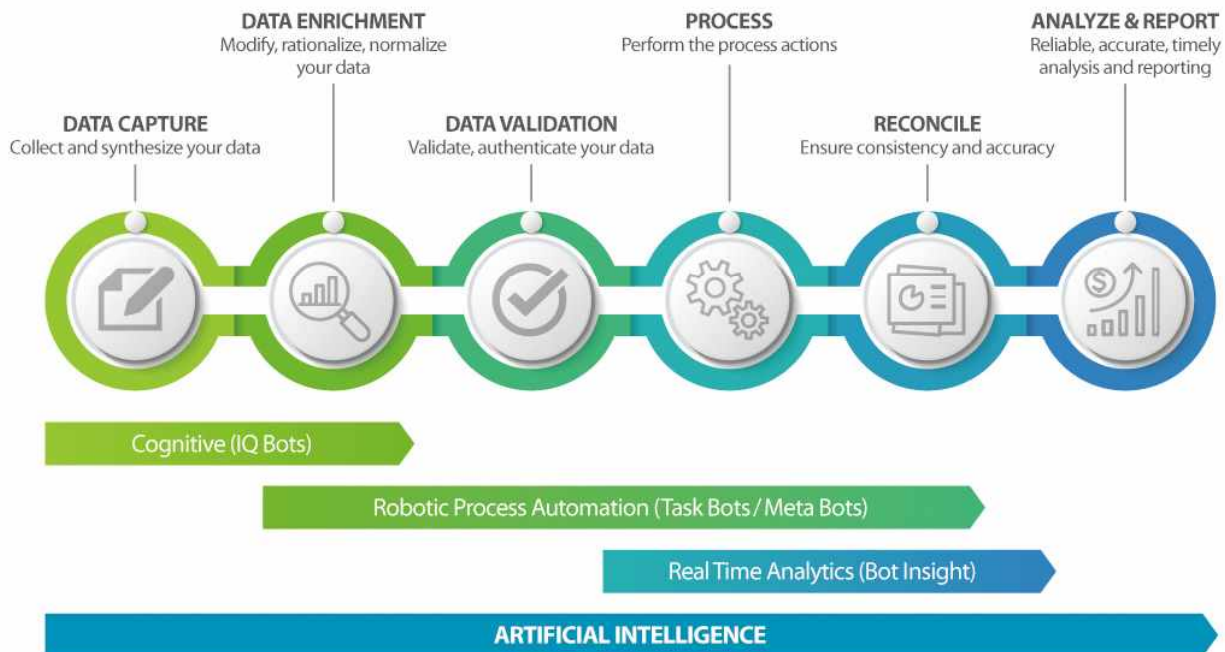
적용 영역

- 일반적인 지식 프로세스 작업에서 인간 노동력을 대규모로 사용하는 곳
- 반복적인 작업을 통해 대량의 트랜잭션 처리 기능을 수행 하는 곳



기대 효과

- 업무 생산성 향상
 - 수분 내에 대량 자동화 업무 실행
 - 24시간 업무 수행
- 휴먼 에러 감소에 따른 업무품질 개선
- 빠른 자동화 구현에 따른 개발 비용 절감
- 전략적/고부가가치 업무 집중



smart solution

KSTEC

Tel : 031-8018-6777 Email : mktg@kstec.co.kr
www.kstec.co.kr

In Partnership with





1599-4855
cheetah@n3ncloud.co.kr



비대면(Untact) AI학습 · 실습 · 개발이 가능한 인공지능플랫폼 '치타'

AI실습플랫폼 치타는 AI실습환경을 인터넷·인트라넷으로 서비스하여 언제 어디서든지 AI실습·연구·개발에 집중할 수 있습니다. 개발 관리자는 컴퓨터 자원과 사용자를 직관적으로 쉽게 관리할 수 있으며, 다수의 연구·개발·실습을 위한 공유환경도 제공합니다.

머신러닝 플랫폼 치타 특징점



Anytime Anywhere AI Platform

인공지능 플랫폼 치타는 학교, 사내, 집, 카페 등 장소에 상관없이, 24시간 접속하여 AI학습·실습·개발이 가능합니다.



1분 안에 구성하는 머신러닝 개발환경

Tensorflow, Keras, PyTorch, Jupyter 등의 설치 없이도, 클릭만으로 1분 안에 머신러닝·딥러닝 환경을 구성합니다.



사용자 환경 (이미지) 저장·관리

인공지능 개발자가 필요에 의해 설치한 라이브러리 등의 개발환경이 저장되어, 인공지능 개발의 연속성을 보장합니다.



손쉬운 사용자 관리

관리자는 프로젝트 그룹의 생성, 사용자의 초대, 등록 등의 관리가 용이하여 사용자를 쉽게 관리할 수 있습니다.



대용량 분산 파일 시스템 지원

Scale out 이 가능한 수백 PB 대용량 분산파일시스템을 지원하여 수백명이 동시 접속해도 병목현상이 없습니다.



용이한 스토리지 관리

관리자는 복잡한 설정없이도 쉽게 스토리지를 운영·관리할 수 있으며, 사용자 중심의 유연한 불륨관리 기능을 제공합니다.



잡스케줄링

Job을 예약할 수 있어서, 연산자원을 효율적으로 운영합니다. Job의 사전검증 기능으로 실행 실패를 최소화합니다.



중앙폴더 공유기능

그룹 유저간의 파일공유가 매우 용이하여, 사용자간 개별 다운로드 및 중복 작업이 필요하지 않습니다.



멀티클러스터링

복수의 GPU 서버들을 하나로 클러스터링하여, 대용량 데이터 처리를 머신러닝·딥러닝으로 쉽고 빠르게 처리할 수 있습니다.



멀티클라우드 제공

AWS, GCP, Tencent 등의 Public Cloud를 이용한 연산 확장이 가능합니다.



NO.1 초등자존감부터 NO.1 중등자신감까지

아이스크림 홈런



초등 스마트 홈러닝 독보적 **1[☆]위**

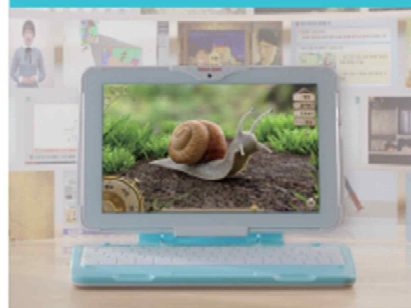
2019 올해의 브랜드 대상 수상! (스마트러닝 부문)

초등학습 상품 시장점유율 1위! (스마트 디바이스 기반)



아이스크림 홈런만의 차별화된 특징점

세계 최대 교육 콘텐츠



우리 아이의 공부 재미와
이해력을 책임지는
AR, VR, 3D 등
300만 개의 미래형 콘텐츠

학교 공부 예습·복습



예비 초부터 중등까지
학습 가능!
전과목 100% 교과 연계로
학교 공부 철저하게 예·복습!

AI 맞춤 학습분석



일 1,000만 건의
학습 데이터를 분석하여
아이의 바른 학습 습관을
만들어주는 홈런 AI생활기록부

[전체 프로그램]

시 간	순 서				
08:00~09:00	한국지능정보시스템학회 이사회				
	SESSION A	SESSION B	SESSION C	SESSION D	SESSION E
09:20~10:50	A1 [학술논문]	B1 [학술논문]	C1 [학술논문]	D1 [학술논문]	E1 [학술논문]
	비즈니스와 산업 좌장: 장영봉 (성균관대)	사용자와 기술 좌장: 양성병 (경희대)	딥러닝 활용 좌장: 채성욱 (호서대)	지능형 기술응용 I 좌장: 이정 (한국외대)	지능형 기술응용 II 좌장: 방영석 (연세대)
10:50~11:10	휴식				
11:10~11:30	개회식				
	개회사: 한국지능정보시스템학회 이정승 조직위원장 환영사: 한국지능정보시스템학회 송용욱 학회장 시상: 감사패 증정, 우수논문상 시상, 인텔리전스대상/우수상 시상				
11:30~12:30	기조강연1: 포스트 코로나 시대 대응을 위한 중소기업 전략 (KISTI 김은선 데이터분석본부장) [현장]				
	기조강연2: 포스트 팬데믹 시대, 교육의 변화와 인공지능의 역할 - 교육 분야가 직면한 이슈와 맞춤형 학습에 대하여 (조용상 아이스크림에듀 대표) [현장]				
12:30~14:00	점심 (open)				
14:00~15:30	A2 [기획세션]	B2 [특별세션]	C2 [특별세션]	D2 [학술논문]	E2 [학술논문]
	디지털 비즈니스 I	인텔리전스 대상 기업세션 I	인공지능 응용 경진대회 I	코로나 19 팬데믹	예측과 의사결정
	좌장: 안현철	좌장: 권영옥	좌장: 이홍주	좌장: 최재영	좌장: 최재원
	(국민대)	(숙명여대)	(가톨릭대)	(아주대)	(순천향대)
15:30~15:50	휴식				
15:50~17:20	A3 [기획세션]	B3 [특별세션]	C3 [특별세션]	D3 [학술논문]	E3 [학술논문]
	디지털 비즈니스 II	인텔리전스 대상 기업세션 II	인공지능 응용 경진대회 II	개인화 기술	콘텐츠 분석
	좌장: 홍준석	좌장: 김남규	좌장: 이홍주	좌장: 천세학	좌장: 박윤주
	(경기대)	(국민대)	(가톨릭대)	(서울과기대)	(서울과기대)
17:20~17:40	폐회식 및 경진대회 시상식				
	사회: 이정 교수(한국외대)				

[세부 프로그램]

A1 [학술논문] 비즈니스와 산업

- A1.1 [Empirical Investigation on Different Aspects of Mobile Game's Success](#) 2
 김정환, 이동원(고려대)
- A1.2 [Exploring Mobility-Based O2O Business Success: A Cross-Country Study on Ride Hailing and Online Food Delivery Services](#) 5
 설재원, 이동원(동국대)
- A1.3 [다분야 서비스 가능 기계학습 엔진 비즈니스 모델 사례 연구](#) 7
 손동성, 황보유정, 이경전(경희대)
- A1.4 [인수합병을 통한 신생기업의 IPO 이후 상태 변화가 산업집중도에 미치는 영향 분석: ICT 산업을 중심으로](#) 9
 장영봉(성균관대), 권영옥(숙명여대)
- A1.5 [인공지능 확산에 따른 산업간 융합 동향](#) 10
 유영재, 윤승주, 김종우(한양대)

B1 [학술논문] 사용자와 기술

- B1.1 [음성쇼핑\(Voice Shopping\)의 혁신저항 및 수용의도에 관한 연구: 가상비서 및 사용자 속성 중심으로](#) 13
 안수호, 김형기, 정두희(한동대)
- B1.2 [로봇바리스타 카페의 서비스스케이프 특성이 고객만족, 즐거움 및 행동의도에 미치는 영향: S-O-R 프레임워크를 기반으로](#) 14
 황주은, 양성병(경희대)
- B1.3 [인플루언서 동반여행의 여행상품 선택속성과 인플루언서 특성이 만족도 및 소비자 반응에 미치는 영향: S-O-R 프레임워크를 기반으로](#) 15
 이은지, 양성병(경희대)
- B1.4 [이차원 고객충성도 세그먼트 기반의 하이브리드 고객이탈예측 방법론](#) 16
 홍승우, 김형수(한성대)
- B1.5 [Users' acceptance toward human-robot interactive service robots : based on YouTube reviews](#) 19
 여방, 장수평, 최재원(순천향대)

C1 [학술논문] 딥러닝 활용

- C1.1 [RNN을 이용한 스마트 홈 환경의 난방 제어에 대한 연구](#) 21
 김부건, 김경옥(서울과기대)
- C1.2 [전문성 이식을 통한 딥러닝 기반 전문 이미지 해석 방법론](#) 22
 김남규, 김태진(국민대)

C1.3	사전 활성화 어텐션 모듈 기반 유방암 이미지 분류 방안	23
	이홍주, 양석우(가톨릭대)	
C1.4	Algorithms vs. Data? LSTM model based Demand Forecasting Using External Weather	24
	박성혁, 김기태(KAIST)	
C1.5	감정 이미지 분류를 위한 CNN과 K-means RGB cluster 앙상블 기법	26
	이홍주, 한기웅, 이정현(가톨릭대)	

D1 [\[학술논문\] 지능형 기술응용 I](#)

D1.1	유전 알고리즘(GA)을 이용한 회계부정 탐지모형(SVM) 변수 최적화에 대한 연구	28
	신경식, 조수현(이화여대)	
D1.2	제로에너지타운 에너지클라우드 데이터 보안 기능 고도화 연구	29
	조충호, 전민재, 광재혁, 이대국, 김민성, 김이강, 윤석호(고려대)	
D1.3	엔진 데이터 구분 방법론을 통한 실시간 실린더 최대압력 및 IMEP 최적화 ELM 예측 비교	31
	이병석(UNIST), 성노윤, 남기환(KAIST)	
D1.4	국민 참여 방식의 SNS 기사단 데이터베이스시스템 구축방안	33
	정기혁(고려대), 한울(상명대)	

E1 [\[학술논문\] 지능형 기술응용 II](#)

E1.1	Siamese Attention LSTM Network를 이용한 문장 ↔ 관련법령 비교 시스템 구축 방안 연구	35
	정지인, 김민태, 김우주(연세대)	
E1.2	계량 측정을 위한 문자 인식 기술 구조	37
	이교혁, 김우주(연세대)	
E1.3	언어 모델링을 활용한 도메인 특화 신경망 기계 번역	38
	이규복, 김우주(연세대)	
E1.4	정보보호 대책의 성능을 고려한 투자 포트폴리오의 게임 이론적 최적화	41
	이상훈, 김태성(충북대)	
E1.5	랜덤성에 기반한 인공지능 오목의 구현	42
	남건민, 천세학(서울과기대)	

A2 [\[기획세션\] 디지털 비즈니스 I](#)

A2.1	Deep Learning 기반 지능형 CCTV	45
	조명돌 - (주)마크애니 최고경영관리본부장	
A2.2	포스트 팬데믹 시대의 AI학습의 중요성	46
	박서린 - (주)N3N CLOUD 매니저	

B2 [특별세션] 인텔리전스 대상 기업세션 I

- B2.1 [uTH 2.0 - 지능형 데이터 기반 디지털무역물류 플랫폼](#) 48
김채미 - (주)한국무역정보통신 디지털무역본부장
- B2.2 [인공지능 챗봇 소개 및 활용 사례](#) 49
장정훈 - (주)와이즈넷 성장기술연구소장
- B2.3 [Hyper Inference - 실사수준 렌더링 활용 객체탐지 학습 솔루션 \(기아자동차 사례중심\)](#) 50
조만희 - (주)메가존 구글클라우드 테크팀
- B2.4 [빅데이터 기반 글로벌 케이팝 플랫폼 "후즈팬"](#) 51
곽영호 - (주)한터글로벌 대표이사

C2 [특별세션] 인공지능 응용 경진대회 I

- C2.1 [시계열 클러스터링을 이용한 이러닝 학습자 유형 변화 분석](#) 53
임종언, 김민지, 남윤주, 류예나, 이채연(국민대)
- C2.2 [중첩적 클러스터링을 활용한 군집별 서비스 제안](#) 54
김민수, 김효중, 신동훈(연세대)
- C2.3 [학습행동에 따른 학생 군집별 중도해지예측모델](#) 55
이재민, 이호진, 김재우, Janine Anne Laddaran, Liu Zih Syuan(연세대)

D2 [학술논문] 코로나 19 팬데믹

- D2.1 [오토인코더 기반 특징추출을 통한 국가별 COVID-19 일일 신규 확진자 수 예측 모델링](#) 57
홍성원(충북대)
- D2.2 [SIR 모형과 LSTM 기반의 COVID-19 확진자 예측 비교 및 분석](#) 59
고석주, 김진오, 김지수, 김희원, 박효상, 김지혜, 옥혜원, 문혜원(경북대), 정명훈(구글 코리아)
- D2.3 [코로나19 접촉 추적 앱의 스트레스 요인이 그 유용성에 미치는 효과](#) 65
주재훈, 한위밍(동국대)
- D2.4 [의료 AI 시스템 개발을 위한 AI 윤리의 기술적 적용과 '설명가능한 AI\(XAI\)' 구현 전략](#) 66
안종훈(인공지능콘텐츠LAB)

E2 [학술논문] 예측과 의사결정

- E2.1 [강화학습 주식투자 수익률 비교에 대한 연구](#) 69
이현재, 이홍주(가톨릭대)
- E2.2 [LSTM분류의 KOSPI 지수 예측에의 적용](#) 71
이재철, 천세학(서울과기대)

E2.3	Factors influencing Investment Decision of Individual Investors in Social Impact Investment using Crowdfunding platform	74
	Pham Duong Thuy Vy, Tran Hung Chuong, 최재원(순천향대)	
E2.4	재무 정보 기반의 기계 학습을 활용한 투자 의사결정	75
	김종우, 전소영, 김동성, 전상경(KAIST)	

A3 [\[기획세션\] 디지털 비즈니스 II](#)

A3.1	RPA 적용 사례 및 AI	78
	이우향 - KSTEC 전문위원	
A3.2	인공지능으로 인한 금융비즈니스 탈바꿈	79
	김업상 - 투이컨설팅 상무	

B3 [\[특별세션\] 인텔리전스 대상 기업세션 II](#)

B3.1	머신 러닝 기반의 자동차 파손 및 심도 인식 모델	81
	김태윤 - (주)에자일소다 AI알고리즘연구팀장	
B3.2	하이퍼오토메이션 실현을 위한 인공지능 기술과 AI InspectorOne	82
	김계관 - (주)그리드윈 대표이사	
B3.3	Redesigning Digital Banking with AI	83
	강정석 - (주)에이젠글로벌 대표이사	
B3.4	집에서 하는 셀프구강검진 '이아포'	84
	양희영 - (주)큐티티 선임연구원	
B3.5	인공지능 QA 자동화 서비스 HBsmith	85
	한종원 - (주)에이치비스미스 대표이사	
B3.6	노타의 Vision AI 모델 자동 압축 기술	86
	이기환 - (주)노타 이사	

C3 [\[특별세션\] 인공지능 응용 경진대회 II](#)

C3.1	SOM-K모델을 이용한 과목별 학습전략 제시(자기조절학습 변인을 중심으로)	88
	원종관, 장영진, 김수민, 김민수, 울가 체르냐예바(부산대)	
C3.2	K-Means와 XGBoost를 사용한 이러닝 플랫폼 이용자 이탈방지 전략 연구	89
	김태영, 조예린, 이세한, 김소연, 윤소정, 이세한(연세대)	

D3 [\[학술논문\] 개인화 기술](#)

D3.1	추천의 다양성을 고려한 그래프 기반 추천시스템	91
	나혜연(UNIST), 남기환(KAIST)	

D3.2	AI 대 전문가의 추천 효율성 비교 실험 연구: 이커머스 내 패션 스타일 추천을 중심으로 93
	김주영, 남기환(KAIST)

E3 [\[학술논문\] 콘텐츠 분석](#)

E3.1	댓글 분석을 통해 제한된 콘텐츠와 제한되지 않은 콘텐츠를 분류---Youtube으로 중심 96
	장수평, 여방, 최재원(순천향대)
E3.2	악성 댓글 분류 시스템 모니터링 연구: 네이버 클린봇 분석 97
	유지웅, 황보유정, 손동성, 이경전(경희대)
E3.3	A Study on the User Perception in ASMR Marketing Content through Social Media Text-Mining: ASMRtist vs Normal Influencer 99
	최재원, Tran Hung Chuong, Pham Duong Thuy Vy(순천향대)
E3.4	Dynamic Topic Model을 활용한 미래 시점 특허 문서 생성에 관한 연구 100
	김우주, 양낙영, 김영, 양재원, 김동욱, 임우담, 이태욱(연세대)
E3.5	BERT기반의 Context Sentence정보를 활용한 계약서 조항 분류 연구 102
	전명주, 김우주(연세대), 임영익, 홍기현, 정병택(인텔리콘 연구소)

A1 [학술논문] 비즈니스와 산업

A1.1 Empirical Investigation on Different Aspects of Mobile Game's Success

이동원
고려대학교 경영대학
mislee@korea.ac.kr

김정환
고려대학교 경영대학
jhkimn124@korea.ac.kr

Abstract – The aim of this study is to develop and validate different aspects of mobile game's success. The dataset of the study consists of 852 free mobile games. Data for the study were collected from the AppAnnie and AppFigure, the mobile data and analytics platforms. The dependent variables of the research are the sales, number of downloads and usage penetration of mobile games, while the main explanatory variables are Addictive, Competition, and Community, which are specific characteristics of mobile game genres. The results of our study indicate that there are positive and significant relationships between Competition and Community, and mobile game sales and usage penetration.

Key Terms – Community, Competition, Game Addiction, In-App-Purchase, Mobile Game

I. Introduction

In various mobile application (app) categories, mobile game is one of the largest and fastest growing ones. Mobile gaming remains the largest segment in 2019, growing 10.2% year on year to \$68.5 billion. Of this, \$54.9 billion will come from mobile games (Newzoo 2019). Recent studies have analyzed a variety of perspectives which could affect the revenue and number of downloads of mobile apps. For example, Krishnan et al. (2019) and Finkelstein et al. (2017) find that ratings could affect positively to the number of mobile app downloads. Ratings have significant effects on the sales of mobile apps (Liu et al., 2015). Lee et al (2014) find that app-level properties such as free price, user review score, and initial rank are more highly associated with survival probability in the higher rank charts. Also, app enjoyment, word-of-mouth (WOM) about mobile app, and app trialability significantly affect the monetary value of the apps (Kim et al., 2016).

However, only few of these studies have investigated the relationship between specific characteristics of mobile game genres—such as Community, Competition, and Addiction, and diverse perspectives of mobile game success—the number of downloads, usage penetration, and sales. In particular, several studies only focus on usage penetration which could be one of most important mobile game indicators, because of the ease of installation and deletion of mobile apps (AppAnnie 2019). Therefore, this study aims to understand the relationship between unprecedented explanatory variables, such as Community, Competition,

Addiction, and dependent variables; Sales, Downloads, and Usage Penetration. In this regard, this study investigates the following hypotheses:

H1. Game Addiction positively affects mobile game success indicators.

H2. Community positively affects mobile game success indicators.

H3. Competition positively affects mobile game success indicators.

Data for the study were collected from the AppAnnie and AppFigure, the mobile data and analytics platforms. Total of 852 distincts games were collected and all of samples were free-to-play. The app data consists of top 1,000 apps in 2018.

In line with the hypotheses of the study, research data was evaluated using Pearson correlation and multiple regression analysis. The relationships between the variables were determined using Pearson correlation analysis. After the Pearson correlation analysis, it was evaluated whether the independent variables significantly influence the dependent variable, by the use of regression analysis.

The results show that Competition and Community have positive relationship with mobile game sales and usage penetration while none of independent variables affect downloads. Also Addictive has no significant relationship with dependents variables.

II. Figures and Tables

<Table 1> Model Variables

Variable	Description
Dependent Variables	
Sales	Monthly sales for a game
Downloads	Monthly downloads for a game
Usage Penetration	The percentage of users who used the app
Explanatory variables	
Addictive	Whether the game genre is one of Action, Card, Casino, RPG, and Strategy (0 = No, 1 = Yes)
Competition	Whether there is competition contents for a game (0 = No, 1 = Yes)
Community	Whether there is user community for a game (0 = No, 1 = Yes)
Update_After_2018	Number of game updates after 2018
Update_Release	Number of game updates after release date

Update_Version	Number of version update	Num_Month_Rank	0.126(0.001)***	-0.078(0.040)**	0.346(0.000)***
Update_Major	Number of major update	Star_Rating	-0.031	0.030	-0.019
Update_Minor	Number of minor update	Ratings	0.016	0.189(0.000)**	0.051
Publisher_Diff	Whether the developer is different with publisher for a game (0 = No, 1 = Yes)	DateDif_month	-0.007	0.071(0.074)*	-0.161
Publisher_minor_major	Whether the publisher is major company	Size	0.001	0.018	0.033
Subsidiary Company	Whether there is parent company for a game developer (0 = No, 1 = Yes)	Rated	0.037	-0.040	0.025
HQ_Country	Whether the developer's HQ is located overseas (0 = No, 1 = Yes)	Num_Languages	-0.013	-0.010	-0.015
IAP_Total_Number	Number of in-app purchase products for a game	Sales_Slot1	-0.016	-0.057	0.167(0.000)**
IAP_Type	Number of in-app purchase products type for a game	Sales_Slot2	-0.035	0.017	0.522(0.000)**
Num_Month_Rank	Number of nomination in monthly top 1000 sales and downloads ranking	Downloads_Slot1	0.053	-0.039	0.03(0.086)*
Star_Rating	User rating score of the game (out of 5)	Downloads_Slot2	0.111(0.002)**	0.015	-0.018
Ratings	Number of user ratings				
DateDif_month	Number of months after game released				
Size	Size of game (MB)				
Rated	Age and content ratings				
Num_Languages	Number of supported languages				
Sales_Slot1	Whether a game locates in Top Grossing Games chart slot 1 (0 = No, 1 = Yes)				
Sales_Slot2	Whether a game locates in Top Grossing Games chart slot 2 (0 = No, 1 = Yes)				
Downloads_Slot1	Whether a game locates in Top Games chart slot 1 (0 = No, 1 = Yes)				
Downloads_Slot2	Whether a game locates in Top Games chart slot 2 (0 = No, 1 = Yes)				

<Table 2> Correlation Analysis Results

III. References

	Addictive	Competition	Community
Addictive			
Competition	-0.043		
Community	0.090**	0.153**	

<Table 3> Multiple Regression Analysis Results

Independent Variables	Downloads	Usage	Sales
Usage Penetration			0.018
In_Average_Downloads		0.000	-0.003
Addictive	-0.005	0.052	0.000
Competition	0.044	0.243(0.000)**	0.034(0.064)**
Community	-0.008	0.079(0.064)*	0.131(0.000)**
Uptdata_After_2018	0.042	-0.025	-0.001
Update_Release	0.026	-0.012	-0.020
Update_Version	0.038	0.038	-0.015
Update_Major	-0.029	0.026	0.001
Update_Minor	-0.035	0.031	-0.007
Publisher_Diff	-0.048	0.041	-0.002
Publisher_minor_major	-0.063(0.096)*	-0.051	0.043(0.021)**
Subsidiary Company	0.061(0.094)*	0.008	0.019
HQ_Country	-0.026	-0.042	-0.010
IAP_Total_Number	0.020	-0.063	-0.002
IAP_Type	0.031	0.029	0.027

Finkelstein, A., Harman, M., Jia, Y., Martin, W., Sarro, F., & Zhang, Y. (2017). Investigating the relationship between price, rating, and popularity in the Blackberry World App Store. *Information and Software Technology*, 87, 119–139.

Kim, H.-W., Kankanhalli, A., & Lee, H.-L. (2016). Investigating decision factors in mobile application purchase: A mixed-methods approach. *Information & Management*, 53(6), 727–739.

Krishnan G, Selvam G. (2019) Factors influencing the download of mobile health apps: Content review-led regression analysis. *HEALTH POLICY AND TECHNOLOGY*. 8(4):356–364.

Lee, G., & Raghu, T. S. (2014). Determinants of Mobile Apps' Success: Evidence from the App Store Market. *Journal of Management Information Systems*, 31(2), 133.

Liu, J., Kauffman, R.J., Man, D., (2015). Competition, cooperation, and regulation: understanding the evolution of the mobile payments technology ecosystem. *Electronic Commerce Research and Applications* 14 (5), 372–391.

Newzoo., (2019). The Global Games Market Will Generate \$152.1 Billion in 2019 as the U.S.

Overtakes China as the Biggest Market. Retrieved June 18, 2019, from
<<https://newzoo.com/insights/articles/the-global-games-market-will-generate-152-1-billion-in-2019-as-the-u-s-overtakes-china-as-the-biggest-market/>>.

A1.2 Exploring Mobility-Based O2O Business Success: A Cross-Country Study on Ride Hailing and Online Food Delivery Services

설재원
고려대학교 경영대학
jaewonseol@gmail.com

이동원
고려대학교 경영대학
mislee@korea.ac.kr

Abstract – Cross-country studies on the success of e-commerce services have revealed that IT infrastructure, economic factors, and sociopolitical factors have profound effects on e-commerce success. This study dug into specifically mobility-based online to offline (O2O) services such as ride-hailing and online food delivery services across more than 130 countries. More specifically, we examine three different categories of influencing factors: mobile infrastructure, economic factors, and empiric risk factors, which could have significant effects on mobility-based O2O business performance. The results of our study indicate that mobile infrastructure plays important roles not only mobile phone penetration but also digital payment capabilities. Moreover, empiric risk factors also negatively affect ride-hailing and online food delivery services' performance. This study is one of the first cross country studies that brings traditional influencing factors on e-commerce to bear mobility-based O2O business performance.

Key Terms – Cross-country study, Empiric risk, Online food delivery, Online to Offline(O2O), Ride-hailing

I. Introduction

Online to offline, commonly abbreviated to O2O, is a channel that entices consumers to purchase offline services or goods through digital platforms, and it has been welcome and grown fast in almost every country with technologically advanced infrastructure (Javalgi and Ramsey, 2001). There are many different types of O2O services such as accommodation, fitness, medical, commerce, delivery, mobility, beauty, etc. In this research, we examine two popular O2O services, ride-hailing and online delivery services closely related to mobility and transportation. Many cross-country studies have investigated relationships among IT infrastructure, physical infrastructure, social/cultural infrastructure, and e-commerce adoption and performance. However, none of the studies have investigated the success of mobility-based O2O businesses globally. Therefore, we investigate following hypotheses:

H1. Mobile and digital payment capabilities are positively related to mobility-based O2O business performance.

H2. The level of economic freedom is positively related to the mobility-based O2O business

performance.

H3. The level of risk of accident or crime is negatively related to the mobility-based O2O business performance.

In this study, we did multiple regression to find the result. The result shows that mobile infrastructure along with digital payment has positive relationship with both ride-hailing and online food delivery services. Interesting point of the finding is that payment channel is more important than mobile phone penetration. In addition, economic factors have positive relationship, while empiric risk shows negative impact.

II. Figures and Tables

<Table 1> Model Variables

Variable	Description
Dependent Variables	
Ride-hailing user penetration	User penetration rate in the Ride-hailing segment in 2019
Online food delivery user penetration	User penetration rate in the Online food delivery segment in 2019
Independent variables	
Mobile_Phone_Subscription 2018	Subscriptions to a public mobile telephone service using cellular technology, including the number of pre-paid SIM cards active during the past three months, expressed as the number of mobile telephone subscriptions per 100 inhabitants in 2018.
MobilePOS_user_penetration 2019	The Mobile POS Payments segment includes transactions at point-of-sale that are processed via smartphone applications so-called "mobile wallets". Payment transactions with physical debit or credit cards at contactless terminals and mobile POS systems in 2019.
DigitalPymnt_user_penetration 2019	The Digital Commerce segment covers all consumer transactions made via the Internet which are directly related to online shopping for products and services. Online transactions can be made via various payment methods (credit cards, direct debit, invoice, or online payment providers, such as PayPal and AliPay).

Economic_Freedom 2019	To measure economic freedom based on 12 quantitative and qualitative factors, grouped into four broad categories, or pillars, of economic freedom: Rule of Law (property rights, government integrity, judicial effectiveness) Government Size (government spending, tax burden, fiscal health) Regulatory Efficiency (business freedom, labor freedom, monetary freedom) Open Markets (trade freedom, investment freedom, financial freedom) Each of the twelve economic freedoms within these categories is graded on a scale of 0 to 100. A country's overall score is derived by averaging these twelve economic freedoms, with equal weight being given to each.
Population_density 2018	Population density is midyear population divided by land area in square kilometers.
Pump_price_for_gasoline 2016	Fuel prices refer to the pump prices of the most widely sold grade of gasoline in 2016. Prices have been converted from the local currency to U.S. dollars.
Crime_rate 2019	Reported crime rate per 100,000 people in 2019. Crime rate is calculated by dividing the number of reported crimes by the total population, and then the result is multiplied by 100,000
Daily_alcohol_intake 2019	Average Daily Intake in milliliters of Pure Alcohol. Pure alcohol consumption among persons (age 15+, drinkers only).

<Table 2> Correlation Analysis Results

	Mobile_Phone_Subscription_2018	Mobile_POS_user_penetration_2019	Digital_payment_user_penetration_2019	Economic_Freedom_Index_2019	Population_density_2018	Pump_price_for_gasoline_2016	Crime_rate_2019	Daily_alcohol_intake_2019
Mobile_Phone_Subscription_2018	1	0.280	0.597	0.458	0.300	0.002	-0.357	-0.018
Mobile_POS_user_penetration_2019	0.280	1	0.307	0.310	0.391	0.256	-0.216	-0.139
Digital_payment_user_penetration_2019	0.597	0.307	1	0.526	0.123	0.250	-0.502	-0.038
Economic_Freedom_Index_2019	0.458	0.310	0.526	1	0.336	0.307	-0.374	-0.095
Population_density_2018	0.300	0.391	0.123	0.336	1	0.204	-0.203	-0.174
Pump_price_for_gasoline_2016	0.002	0.256	0.250	0.307	0.204	1	-0.243	-0.050
Crime_rate_2019	-0.357	-0.216	-0.502	-0.374	-0.203	-0.243	1	0.130
Daily_alcohol_intake_2019	-0.018	-0.139	-0.038	-0.095	-0.174	-0.050	0.130	1

		Ride-hailing user penetration		Online food delivery user penetration	
		Beta	Sig	Beta	Sig
IT	Mobile_Phone_Subscription	0.234 ***		0.030	
	Mobile_POS_user_penetration	0.303 ***		0.295 ***	
	Digital_payment_user_penetration	0.427 ***		0.674 ***	
Economic	Economic_Freedom	0.412 ***		0.532 ***	
	Population_density	0.319 ***		0.101	
	Pump_price_for_gasoline	-0.173 **		0.144 *	
Empiric_risk	Crime_rate	-0.321 ***		-0.418 ***	
	Daily_alcohol_intake	-0.245 ***		-0.196 **	

"A Growth Theory Perspective on E-Commerce Growth in India: An Exploratory Study," *Electronic Research and Applications* 237-259.

Javalgi, R., and Ramsey, R. 2001. "Strategic Implications of E-Commerce as an Alternative Distribution System," *International Journal of Electronic Commerce* (6:12), p. 10.

Kraemer, K. L., Ganley, D., and Dewett, T. 2007. "Across the Digital Divide: A Cross-National Multi-Technology Analysis of Determinants of It Penetration," *Journal of Management Information Systems* (24:1), pp. 443-452.

Kshetri, N. 2007. "Barriers to E-Commerce Adoption in Developing Countries: A Case Study," *Journal of Electronic Commerce Research and Applications* (7:1), pp. 22-35.

Li, P., and Xie, W. 2012. "A Strategic Framework for Determining E-Commerce Adoption in China," *Journal of Technology Management* (7:1), pp. 22-35.

A1.3 다분야 서비스 가능 기계학습 엔진 비즈니스 모델 사례 연구

이경전
경희대학교 경영대학
klee@khu.ac.kr

손동성
경희대학교 이과대학
chdtjd92@khu.ac.kr

황보유정
경희대학교 이과대학
hwangbo@khu.ac.kr

Abstract - 인터넷 기술의 발전이 플랫폼 비즈니스 모델을 발전시킨 것과 같이, 인공지능 기술, 특히 기계학습 기술이 과연 프로젝트형, 솔루션 모델을 넘어서 플랫폼 비즈니스 모델과 같은 새로운 비즈니스 모델을 발전시킬 것인가에 대한 산업과 학계의 관심이 늘어나고 있다. 본 논문에서는 에듀테크 기업 위이드(Riio) 사례를 통해 기계학습 엔진에 기반한 다분야 서비스 비즈니스 모델의 가능성을 탐색적으로 연구한다. 이 모델은 네트워크 효과에 기반하지 않는다는 면에서 인터넷 기반의 플랫폼 비즈니스 모델과 차별성을 가지며, 데이터 효과가 발생한다는 점에서 인공지능 기반의 비즈니스 모델의 중요한 특성 하나를 도출하며, 일반 인공지능(AGI: Artificial General Intelligence)이 추구하는 것이 너무 이상적이고 기술중심적인 상황에 놓여 있는 것에 비해, 기계학습 엔진에 기반하여 다분야에 서비스를 가능하게 하면서 수익을 추구하고 있다는 면에서, 세계적으로 더 유명한 인공지능 기업들 보다 더 우월한 비즈니스 모델을 보여주고 있다.

Key Terms - AI, Case Study, Cross Domain Service, AI Business Model, Machine Learning, Machine Learning Engine

사사: 2017년 대한민국 교육부와 한국연구재단의 지원을 받아 수행된 연구임
(NRF-2017S1A5B8059804)

I. 서론

인터넷의 급격한 발전에 따라 구글(Google), 아마존(Amazon), 우버(Uber), 페이스북(Facebook)과 같은 새로운 플랫폼 비즈니스 모델이 등장하였다. 이들은 직접 서비스나 제품을 제공하지 않고, 생산자와 사용자를 연결하여 거래가 발생하도록 함으로써 수익을 창출한다. 한편, 2016년 인간보다 바둑을 더 잘 두는 인공지능 알고리즘이 등장하여 많은 사람들에게 충격과 기대를 주었다(Silver et al. 2016). 인터넷의 등장이 다양한 플랫폼 비즈니스 모델을 만들어낸 것과 같이, 인공지능 기술을 기반으로 하는 어떤 새로운 비즈니스 모델들이 탄생될 것인가? 탄생된다면, 무슨 기술로, 어떤 형태로 만들어질 것인가? 이러한 질문은 경영학, 비즈니스 모델, AI 분야 모두의 관심사이다. 위이드는 2014년 설립된 시험문제 형태의 교육 콘텐츠 서비스를 제공하는 에듀테크 회사이다. 객관식 시험 영역에서 특화된 협동적 필터링, 딥러닝, 강화학습 기술에 기반한 엔진을 개발하여 1:1 맞춤형 서비스 산타토이를 만들었으며, 일본에 TOEIC, 베트남에 SAT 등 Test-prep 시장에 진출하고 있다. 사용자에게 부족한 영역을 추천하여 단시간에 성적을 올릴 수 있게 하여 인기를 끌고 있으며, 지불의사(WTP: Willingness to Pay)가 높은 TOEIC 분야를 먼저 공략했다는 점에서 비즈니스 모델의 우

수 사례로 판단된다. 본 연구는 Riio의 비즈니스 모델을 분석하는 과정에서 인공지능 비즈니스 모델의 특성을 탐색한다.

II. 네트워크 효과 VS. 데이터 효과

1990년대부터 본격화된 인터넷의 대중화와 상업화는 새로운 비즈니스 모델의 등장을 예고했다. 인터넷은 생산비용과 거래비용, 조직화 비용을 동시에 줄여서, 기존의 생산, 유통 모델을 해체시킬 수 있었고, 이러한 상황에서 비즈니스모델의 창출과 설계가 매우 중요해졌다. 이러한 인터넷 기반 비즈니스 모델을 이해하고 설명하는 데에 가장 중요한 개념 중의 하나는 네트워크 효과(Network Effects)였다. 네트워크 효과는 같은 그룹의 사용자가 증가함에 따라 해당 그룹의 사용자의 가치가 같이 증가하는 직접 네트워크 효과와 한 그룹의 사용자가 증가함에 따라 다른 그룹 사용자의 가치가 증가하는 간접 네트워크 효과로 나뉜다(Katz & Shapiro, 1985). 인공지능 기반 비즈니스 모델이 인터넷 기반 비즈니스 모델과 다른 점은, 인공지능 기반 비즈니스 모델은 데이터 효과(Data Effect)가 중요하다는 것이다. 데이터 효과란 제품과 서비스, 비즈니스 모델에 있어, 데이터가 쌓임에 따라 개인과 기업 등의 조직의 경쟁력과 가치가 커지는 것을 의미하는 용어로 제안되었다(이경전, 2019). 위이드는 현재 네트워크 효과가 발생하는 것으로 판단되지는 않지만, 데이터 효과는 발생하는 것으로 판단된다. 위이드의 모델은 기계학습, 즉 학습 기계(Learning Machine)가 맞춤형 서비스를 생산하는 Stabell & Fjeldstad (1998)의 가치 배열(Value Configuration) 유형 중 가치사슬(Value Chain)에 가깝다. 기계에 들어가는 원재료, 즉, 시험문제 콘텐츠를 공급받아 이를 기계학습 엔진이 가공하여 서비스로 제공하는 형태이다. 시험문제공급자는 사용자의 사용량에 따라 수익이 증가하는 구조가 아니다. 만약 수익이 사용자의 사용량에 따라 비례하여 지급되는 구조로 변화한다면, 시험 문제 공급 기업들은 사용자가 많은 위이드의 플랫폼에 진입하기를 희망할 것이며, 이는 간접 네트워크 효과를 발생시킬 것이다. 그 이전까지는 위이드를 플랫폼이라 부르는 것은 엄밀한 의미에서 옳지 않은 것으로 판단된다. 결론적으로 위이드의 사례를 통해서 우리가 어떤 인공지능 기반 비즈니스 모델을 플랫폼으로 부를 것인가 말 것인가의 하나의 기준을 얻게 되었는데, 그것은 네트워크 효과이다. 그리고, 인공지능에 기반한 비즈니스 모델의 특징은 그것이 데이터 효과를 보이고 있는가 하는 점일 것이라는 점이 발견되었다.

III. AGI VS. 다분야 서비스 가능 기계학습 엔진

기존의 상당수의 AI 회사는 고객사의 수주를 받는 용역회사 형태가 대다수였다. 고객사로부터 업무와 데이터를 제공받아 데이터 분석을 수행하고, 분석된 데이터를 기반으로 시스템을 커스터마이즈하여 고객사에게 제공하는 프로젝트 기반의

비즈니스이다. 이는 가치 배열 유형에서 Value Shop이다. 즉, 제품을 생산하여 가치를 창출하는 Value Chain도 아니고, 두 종류의 고객을 연결하는 Value Network도 아닌, 고객의 문제를 해결해주는 것으로, 단발성 계약이 많아 폭발적 성장은 어렵다. 그런데, 의료 분야는 다를 수 있다. 뷰노(VUNO)는 특정 병원의 데이터를 활용하여 개발된 AI 시스템을 공동으로 특허를 출원하고, 공동사업을 통해 다른 병원에 판매하는 모델을 개발하여 병원이 수요자이자 공급자가 되는 형태의 플랫폼 모델의 모습을 보여주기도 했는데, 이는 이질적 네트워크와 동질적 네트워크가 결합된 플랫폼의 형태를 띠는 점에서 흥미로우며(정태훈 & 이경전 2000; 이경전 & 전형원 2000), 가치 배열 유형 중 Value Network에 해당한다.

한편, 퀴이드는 하나의 기계학습 엔진을 가지고 여러 분야의 서비스를 제공하는 것이 가능하다. 현재 퀴이드는 외국어 분야(영어)에 집중되어 있지만, 더 나아가 다른 도메인인 수학, 국어와 같이 시험 문제 콘텐츠를 가진 다양한 기업들과 함께 교육 콘텐츠 서비스를 개발하고 수익을 공유하는 인공지능 기반 교육 서비스 회사로 성장이 가능하다. 앞서 설명한 뷰노의 경우는 하나의 기계학습 엔진으로 다른 분야에 서비스를 적용하는 것이 어렵지만, 퀴이드의 경우는 하나의 엔진으로 여러 분야에서 데이터를 축적하여 모델을 만들면, 각 서비스에 바로 적용이 가능하므로 다분야 서비스 가능 기계학습 엔진 비즈니스 모델로 명명할 수 있다. 영국의 딥마인드와 같은 회사가 AGI(Artificial General Intelligence)를 지향하고 있으나 아직 상업적 성과는 내지 못하는 반면, 한국의 퀴이드는 하나의 기계학습 엔진으로 다분야의 서비스를 제공하면서 수익을 창출하고 있다는 점에서 주목할 만하다.

퀴이드의 비즈니스 모델은 디지털 미(Digital Me)라는 관점에서도 볼 수 있다. 디지털 미는 사용자 개인정보를 추적, 관리해주는 서비스로 헬스, 재정, 교육 등 다양한 분야에서 사용자 개인을 관리해주는 통합 서비스다. 어떤 기업의 비즈니스 모델은 하나의 관점에서만 규정되는 것은 아니며, 본 연구는 인공지능 기술에 의해서 탄생하는 새로운 기업들을 통해서 새로운 비즈니스 모델의 특징을 연구한 초기 사례 연구로서 의미가 있다고 하겠다.

참고문헌

David, S., Huang, A., Maddison, C. J., Guez, A., Sifre, L., Driessche, G. V. D., Schrittwieser, J., Antonoglou, I., Panneershelvam, V., Lanctot, M., Dieleman, S., Grewe, D., Nham, J., Kalchbrenner, N., Sutskever, I., Lillicrap, T., Leach, M., Kavukcuoglu, K., Graepel, T. and Hassabis, D., "Mastering the game of Go with deep neural networks and tree search", Nature: 529 (7587): 484~489, 2016.

Katz, M. L. & Shapiro, C., "Network externalities, competition, and compatibility", The American economic review, 1985.

Stabell, C. B. & Fjeldstad, Ø. D., "Configuring value for competitive advantage: on chains, shops, and networks." Strategic management journal 19(5): 413-437, 1998.

이경전, "[전문가 포럼] AI 시대에 맞는 비즈니스

모델 발굴해야", 한국경제, 2019.02.11

이경전, 전형원, "전자시장에서의 네트워킹 비즈니스 모델", 한국경영과학회 추계학술대회 논문집, 2000

정태훈, 이경전, "인터넷 비즈니스 네트워커에 대한 사례 연구: (주)원큐의 비즈니스 모델, 전략, 기술을 중심으로" 한국경영과학회, 2000

A1.4 인수합병을 통한 신생기업의 IPO 이후 상태 변화가 산업집중도에 미치는 영향 분석: ICT 산업을 중심으로

성균관대학교 경영대학*
장영봉 (ybchang01@skku.edu)

&

숙명여자대학교 경영학부**
권영옥 (yokwon@sm.ac.kr)

2020년 6월

* 이 논문은 2019년 대한민국 교육부와 한국연구재단의 인문사회분야 중견연구자 지원사업의 지원을 받아 수행된 연구임(NRF-2019S1A5A2A01036720)

** 교신저자

초록

본 논문에서는 IPO 이후 기업의 상태를 추적하여, IPO 이후 빠른 시기에 기존 기업에 인수합병되는 상황으로 야기된 산업의 집중도 변화에 대하여 실증적으로 규명하였다. 2000년대 이후 IPO를 단행한 기업들의 경우 이전 세대에 IPO한 기업 대비 상대적으로 빠른 시기에 기존 기업에 인수합병된 것으로 나타났다. 하지만 이러한 기업들의 규모, 수익성 및 연구개발비 투자 등은 기타 이유로 시장에서 퇴출된 기업 대비 양호한 것으로 나타났다. 본 연구의 분석 결과에 의하면, 동일산업군 내 IPO 이후 단기간에 인수되는 기업의 사례가 증가함에 따라 관련 산업의 집중도 역시 증가한 것으로 나타났다. IPO 이후 빠른 시기에 인수된 기업의 수가 증가함에 따라 심화된 집중화 정도는 산업 내 지배적 기업이 존재할 경우 더 크게 영향을 미치는 것으로 나타났다. 하지만 지배적 기업의 비중이 산업 집중도에 미치는 정(+)의 효과가 모든 산업 부문에서 나타난 것은 아니며 ICT 산업 부문에 한정된 것으로 추정되었다.

키워드: ICT 기업/산업, 인수합병, 기업공개, 산업 집중도

A1.5 인공지능 확산에 따른 산업간 융합 동향¹⁾

유영재
한양대학교
일반대학원
경영학과

uyeongjae@hanyang.ac.kr

윤승주
한양대학교
일반대학원
비즈니스인포메틱스학과

thingjoo@hanyang.ac.kr

김종우
한양대학교
경영대학
경영학부

kjw@hanyang.ac.kr

Abstract – 오늘 날 인공지능은 산업 전반에 걸쳐 다양한 분야에서 활용되고 있다. 그리고 기술의 수준이 발전함에 따라 기업들은 인공지능을 매개로 한 교류를 통해 새로운 제품 및 서비스 나아가 새로운 산업 분야 창출을 통해 경쟁력을 확보하고 있다. 따라서 본 연구에서는 인공지능과 관련하여 발생되는 기업 및 산업간의 교류를 온라인 뉴스 데이터를 통해 시계열적으로 분석하여 그 흐름과 동향에 대해 확인해보고자 한다. 산업별 기업들의 인공지능 관련 온라인 뉴스 데이터를 수집하여 문장 단위로 분리한 후 형태소 분석을 진행하였다. 그리고 하나의 문장에서 두개 이상의 기업명과 기업간의 관계가 있음을 의미하는 명사 또는 동사가 발견되면 해당 기업을 연결하였다. 이는 기업 및 산업간의 융합을 정성적으로 확인하는 기존 연구 방법과 달리 뉴스 기사를 통해 정량적이고 빠르게 확인할 수 있으며, 미공개 단계가 지난 후에 분석이 가능한 특허 데이터 기반의 연구 방법과 달리 실시간 데이터를 활용하여 현시점에 대한 분석이 가능하다.

I. 서론

산업 간의 기술 융합은 기존과 차별화된 제품 및 서비스 제공과 이를 통한 새로운 산업 분야 창출까지 연쇄적인 영향을 미치며, 기업은 이를 통해 협력하고 혁신하며 경쟁력을 확보하기 때문에 중요한 현상이다(Hacklin, 2007). 산업 융합의 예로 패션 산업과 IT 산업이 융합하여 만들어진 웨어러블 디바이스가 있으며 그 밖에도 스마트 그리드, 스마트 홈 등이 있다. 따라서 본 연구에서는 네트워크 분석을 통해 인공지능 기술을 사용한 기업을 중심으로 산업간의 교류 및 변화 추이를 확인해보고자 한다. 이를 통해 기업 및 산업별 교류의 변화와 교류가 활발히 일어난 산업들과 그렇지 못한 산업들을 확인해 볼 수 있으며, 향후 인공지능 기술을 매개로 한 산업간의 융합 트렌드 예측을 위한 기초 자료로 활용이 가능할 것으로 기대한다.

II. 관련 연구

2.1 산업간의 융합 분석
산업간의 융합을 파악하기 위해 기존 연구에서는 전문가의 의견에 의존하는 정성적인 방법을 기반으로 연구가 진행되어 왔으나 이는 많은 자원과 시간이 투자된다는 한계점이 있다(An et al., 2016). 다른 한편으로 산업간 융합의 신호는 기술적인 결과로부터 확인 가능하여 주로 특허

데이터를 중심으로 산업간의 융합에 대한 분석이 진행이 되었다(Karvonen and Kassi, 2013). 하지만 특허 정보는 일정기간의 미등록기간이 지난 후에 공개되므로 현시점을 기준으로 한 산업간의 융합 분석이 어렵다는 단점이 있다.

2.2 소셜미디어 데이터 분석

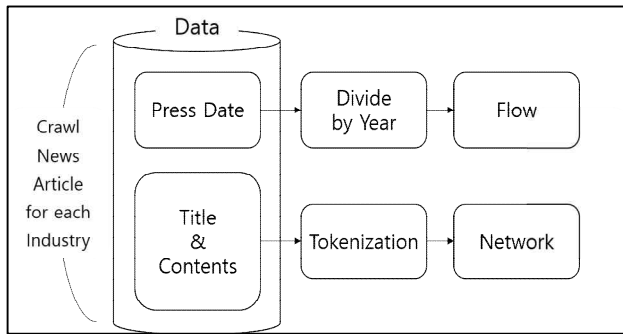
오늘날 소셜미디어는 다양한 기능을 통해 개인 또는 조직의 커뮤니케이션 수단의 역할뿐만 아니라 네트워크 형성에도 기여하며 많은 양의 데이터를 생산하고 있어 기업, 정부 등은 이를 활용하기 위한 방안을 다방면으로 모색하고 있다(Kietzmann et al., 2011; Zeng et al., 2010). 현재 소셜미디어에 대한 연구는 데이터를 수집하여 모니터링, 분석, 요약 및 시각화하는 것이 주를 이루고 있고, 누구나 언제든지 쉽게 실시간으로 이용 및 접근 가능하다는 특징을 갖고 있어 여론 감성 분석, 트렌드 분석 등에 적극 활용되고 있다(Fan and Gordon, 2014; Zeng et al., 2010).

III. 연구방법

본 연구는 <그림 1>과 같은 절차를 통해 진행되었다. GICS(Global Industry Classification Standard) 기준에 따라 산업별로 시가총액 상위 10개 기업을 각 산업의 대표기업으로 선정하였다. 그 후, 국내 뉴스 플랫폼 네이버 뉴스에서 2016년 1월부터 2019년 12월까지 발행된 각 기업의 인공지능 관련 기사를 웹 크롤링(web crawling)을 통해 보도 날짜, 기사 제목, 기사 본문을 수집하였다. 수집된 데이터 중 본 연구를 진행함에 있어 불필요한 불용어, 문장부호, 광고 등을 제거하였다.

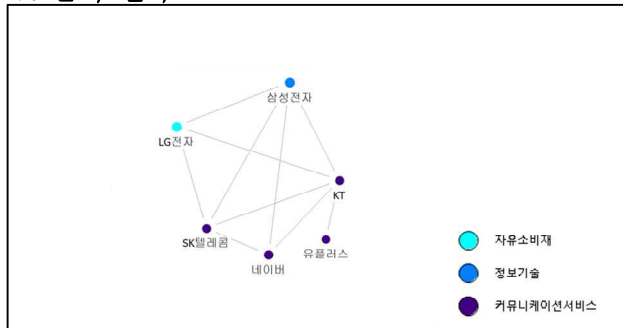
인공지능 발전에 따른 산업간 네트워크를 분석하기에 앞서 각 기사를 문장 단위로 분리하여 형태소 분석을 진행하여 단어의 빈도수를 기준으로 나열하였다. 그 중 상위에 출현한 단어들을 대상으로 융합과 관련된 단어들을 추출하였고, 추출된 단어로는 협력, 공동, 제휴, 협업, 체결, 파트너 등 총 13개가 있다. 그 후, 하나의 문장에서 융합 관련 단어들과 두 개 이상의 기업명이 동시 출현하였을 경우 해당 기업들을 쌍방향으로 연결하였다. 따라서 본 연구의 네트워크 결과는 기업 또는 산업이 인공지능이라는 매개체를 통해서 교류가 일어났다고 볼 수 있다.

* 이 논문 또는 저서는 2017년 대한민국 교육부와 한국연구재단의 지원을 받아 수행된 연구임 (NRF-2017S1A3A2066740).

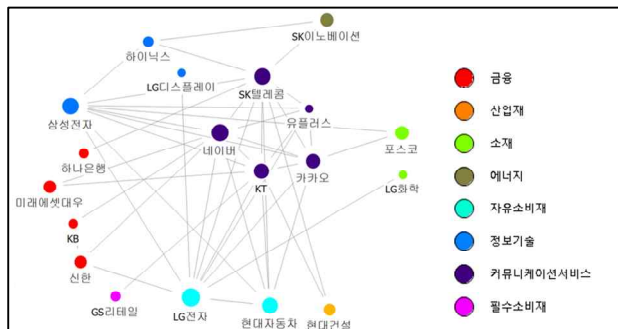


<그림 1> 연구 절차.

IV. 분석 결과



<그림 2> 2016년 연결망 분석 결과.



<그림 5> 2019년 연결망 분석 결과.

2016년에는 3개의 산업(커뮤니케이션서비스, 정보기술, 자유소비재)만이 인공지능을 매개로 하여 융합하고 있었지만, 2019년에는 산업 전반에 걸쳐 융합이 일어나고 있음을 확인할 수 있었다. 커뮤니케이션서비스 산업의 KT와 필수소비재 산업의 GS리테일은 미래형 점포 구축을 위한 협약을 맺은 것을 확인할 수 있었고, 자유소비재 산업의 현대자동차와 커뮤니케이션 산업의 카카오는 카카오 미니의 음성인식 기술을 활용한 차량 모듈 제어 기술을 개발 중인 것을 확인할 수 있었다.

V. 참고문헌

An, J., K.W. Kim, H.Y. Noh, and S.J. Lee, "Identifying Converging Technologies in the ICT Industry: Analysis of Patents Published by Incumbents and Entrants", Journal of the Korean Institute of Industrial Engineers, Vol.42, No.3(2016),

209~221.

Fan, W. and M. D. Gordon, "The Power of Social Media Analytics," Communications of the ACM, Vol.6 No.57(2014), 74~81.

Hacklin, F., "Management of Convergence in Innovation: Strategies and Capabilities for Value Creation Beyond Blurring Industry Boundaries", Springer Science & Business, 2007.

Karvonen, M., T. Kassi, "Patent Citations as a Tool for Analyzing the Early Stages of Convergence", Technological Forecasting and Social Change, Vol.80, No.6(2013), 1094~1107.

Kietzmann, J. H., K. Hermkens, I. P. McCarthy, and B. S. Silvestre, "Social Media? Get Serious? Understanding the Functional Building Blocks of Social Media," Business Horizons, Vol.3, No.54(2011), 241~251.

Zeng, D., H. Chen, R. Lusch, and S. H. Li, "Social Media Analytics and Intelligence," IEEE Intelligent System, Vol.6, No.25(2010), 13~16.

B1 [학술논문] 사용자와 기술

B1.1 음성쇼핑(Voice Shopping)의 혁신저항 및 수용의도에 관한 연구: 가상비서 및 사용자 속성 중심으로

A Study on the Innovation Resistance and Behavioral Intention of Voice Shopping: Focused on Virtual Assistant and User Characteristics

안수호 (Suho Ahn)

소속(국문): 한동대학교 ICT창업학부

소속(영문): School of Global Entrepreneurship and Information Communication Technology, Handong Global University

E-mail: 21600840@handong.edu

김형기 (Hyeongki Kim)

소속(국문): 한동대학교 경영경제학부

소속(영문): School of Business Management and Economics, Handong Global University

E-mail: 21400236@handong.edu

Corresponding author: 정두희 (Doohee Chung)

558, Handong-ro, Heunghae-eup, Buk-gu, Pohang-si, Gyeongsangbuk-do, Republic of Korea

Tel: +82-54-260-1510, Fax: +82-54-260-1149, E-mail: profchung@handong.edu

국문초록

가상비서가 확산됨에 따라, 이를 기반으로 상거래 활동을 하는 음성쇼핑(Voice shopping)이 주목받고 있다. 음성쇼핑은 음성으로 모든 조작을 하는 사용자 친화적 쇼핑 방식이며, 온라인쇼핑에 비해 편의성이 크고 시간적 공간적 제약이 줄어들기 때문에 괄목할 성장이 전망된다. 이 연구에서는 기술수용모델과 혁신저항이론을 결합, 가상비서의 속성과 사용자 속성이 음성쇼핑의 저항과 수용의도에 어떻게 영향을 주는지 분석했다. 음성인식 기술과 비교적 친숙한 20~30대 소비자를 대상으로 설문조사를 했으며, 151명의 유효 응답을 토대로 분석을 진행했다. 정확성, 사회적 실재감, 상호작용성 등 가상비서 속성과 사용자 혁신성 및 사용자 경험 등 사용자 속성을 독립변수로 설정했다. 이러한 변수가 기술수용모델과 혁신저항이론의 주요 변수인 상대적 이점, 인지된 용이성, 인지된 위험 등을 매개로 하여 음성쇼핑에 대한 혁신저항과 수용의도에 미치는 영향을 분석했다. 분석 결과, 혁신저항과 상대적 이점은 수용의도에 긍정적인 영향을 주는 것으로 확인됐으며 상대적 이점과 인지된 용이성, 인지된 위험 모두 혁신저항에 긍정적 영향을 주는 것으로 확인됐다. 한편 가상비서 특성 변수인 정확성은 인지된 위험에만 긍정적 영향을 주었고, 사회적 실재감은 상대적 이점과 인지된 용이성, 상호작용성은 인지된 용이성에 긍정적 영향을 주었다. 사용자 특성 변수인 사용자 경험은 인지된 용이성에만 긍정적 영향을 주었다. 이 연구는 기술수용모델 계열 연구의 단점인 친혁신적편향 성격을 보완하기 위해 혁신저항이론을 결합해 균형 있는 접근으로 분석을 시도했다. 음성쇼핑 플랫폼 및 상거래 서비스를 개발하는 기업들의 서비스 구축에 참고할 수 있는 지침을 제시한다는 점에서 의의를 지닌다.

주제어

음성쇼핑, 가상비서, 혁신저항, 기술수용, 상호작용성.

B1.2 로봇바리스타 카페의 서비스스케이프 특성이 고객만족, 즐거움 및 행동의도에 미치는 영향: S-O-R 프레임워크를 기반으로

The Effects of Robot-Barista Café Servicescape on Customer Satisfaction, Pleasure, and Behavioral Intention: Based on the S-O-R Framework

황주은 (JuEun Hwang)

소속(국문): 경희대학교 경영대학원 경영학과

소속(영문): School of MBA, KyungHee University

E-mail: hyehye0415@gmail.com

양성병 (Sung-Byung Yang)

소속(국문): 경희대학교 경영대학 경영학과

소속(영문): School of Management, KyungHee University

E-mail: sbyang@khu.ac.kr

Corresponding author: Sung-Byung Yang

26 Kyungheedae-ro, Dongdaemun-gu, Seoul, 02447, Korea

Tel: +82-2-961-9548, Fax: +82-2-961-0515, E-mail: sbyang@khu.ac.kr

최근 5G, 인공지능(AI) 등 정보통신기술(ICT)의 발달로 IT와 외식산업이 융합된 다양한 형태의 푸드테크(FoodTech) 기술 및 비즈니스 모델이 등장하고 있다. 이러한 푸드테크 활성화의 예로 외식업계에서 빠르게 적용되고 있는 로봇(서빙 로봇, 셰프 로봇, 바리스타 로봇 등)을 들 수 있는데, 최근 업계의 많은 관심에 비해, 관련 연구는 전무한 상황이다. 이에 본 연구는 S-O-R(Stimulus-Organism-Response) 프레임워크 및 푸드테크 및 로봇 관련 선행연구를 바탕으로 로봇바리스타 카페 맥락에서 서비스스케이프 특성이 사용자의 즐거움과 만족, 나아가 행동의도에 미치는 영향에 대해 실증해 보고자 한다. 또한, 이러한 영향이 카페 라이프스타일(편의추구형, 유행선도형)에 따라 어떻게 달라지는지에 대해 추가로 검증해 보고자 한다. 실증 분석을 위해 수집된 자료의 가설 검증 결과, 로봇바리스타 카페의 기술적 요인인 일관성과 효율성, 디자인적 요인인 혁신성의 경우, 즐거움과 만족 모두에 유의한 영향을 미친 반면, 디자인적 요인의 심미성은 즐거움에만 유의한 영향을 미치는 것으로 나타났다. 즐거움 역시 만족에 유의한 영향을 미치는 것으로 나타났으며, 즐거움과 만족 모두 재방문의도와 구전의도에 유의한 영향을 미치는 것으로 확인되었다. 또한, 카페 라이프스타일(편의추구형 및 유행선도형)의 조절효과 분석 결과, 효율성이 즐거움 및 만족에 미치는 영향은 편의추구형 그룹에서 더 큰 것으로 조사되었다. 본 연구는 S-O-R 프레임워크를 활용하여 로봇바리스타 카페의 맥락에서 서비스스케이프 특성이 소비자의 즐거움과 만족 및 행동의도(재방문의도, 구전의도)에 미치는 영향을 실증한 최초의 연구라는 점과 카페 라이프스타일의 조절효과를 추가로 검증했다는 점에서 연구의 이론적 의의를 찾을 수 있다. 본 연구결과를 통해, 로봇을 활용한 외식업 마케팅 담당자, 창업자 및 운영자에게 유용한 정보를 제공하고, 구체적으로 로봇바리스타 카페의 어떠한 서비스스케이프적 특성을 좀더 강조해야 할지에 대한 가이드라인을 제시할 수 있을 것으로 기대한다.

주제어

로봇바리스타 카페, 서비스스케이프 특성, 고객만족, 즐거움, 재방문의도, 구전의도, S-O-R 프레임워크

B1.3 인플루언서 동반여행의 여행상품 선택속성과 인플루언서 특성이 만족도 및 소비자 반응에 미치는 영향: S-O-R 프레임워크를 기반으로

The Impacts of Selection Attributes of Travel Product and Influencer Characteristics in the Travel with Influencer on Customer Satisfaction and Responses: Based on the S-O-R Framework

이은지 (Eunji Lee)

소속(국문): 경희대학교 경영대학원 경영학과

소속(영문): School of Management, KyungHee University

E-mail: eunji9200@gmail.com

양성병 (Sung-Byung Yang)

소속(국문): 경희대학교 경영대학 경영학과

소속(영문): School of Management, KyungHee University

E-mail: sbyang@khu.ac.kr

Corresponding author: Sung-Byung Yang

26 Kyungheedae-ro, Dongdaemun-gu, Seoul, 02447, Korea

Tel: +82-2-961-9548, Fax: +82-2-961-0515, E-mail: sbyang@khu.ac.kr

국문초록

소셜 네트워크 서비스(Social Network Service)가 대중화되면서 타인에게 큰 영향력을 미치는 인플루언서가 등장하였고, 많은 분야에서 그 영향력이 커지고 있다. 최근 여행업계에서는 전문가, 연예인, SNS 스타 등과 같은 인플루언서들을 동반하는 여행상품들이 소비자들의 큰 인기를 끌고 있다. 그럼에도 불구하고 인플루언서 동반여행 관련 연구는 아직 충분히 이루어지지 않고 있다. 이에 본 연구는 인플루언서 동반여행의 여행상품 선택속성과 인플루언서의 특성이 여행의 만족도 및 반응에 미치는 영향력을 실증해 보고자 한다. 또한, 조절변수로 인플루언서 유형(전문인, 유명인)을 고려하여, 인플루언서 유형에 따라 여행상품 선택속성 및 인플루언서 특성이 만족도에 미치는 영향이 어떻게 달라지는지에 대해서도 살펴보고자 한다. 연구문제의 실증적 규명을 위해 자료 수집 후 가설검증 결과, 여행상품 선택속성 중 여행사서비스와 현지프로그램이, 인플루언서 특성 중 매력성, 전문성 및 공감성이 만족도에 유의한 영향을 미치고, 만족도 또한 브랜드 이미지 및 추천의도에 유의한 영향을 미치는 것으로 나타났다. 또한, 인플루언서 유형(전문인, 유명인)에 따른 조절효과 검증 결과, 매력성이 만족도에 미치는 영향은 유명인과 함께 하는 인플루언서 동반여행의 경우에 더 큰 것으로 조사되었다. 기존 선행연구는 인플루언서를 이용한 홍보 및 마케팅 관점의 연구가 대부분이었다. 본 연구는 한 발 더 나아가, 인플루언서가 직접 상품에 참여하는 인플루언서 동반여행 맥락에서 최초로 살펴본 연구이다. 본 연구결과를 통해 인플루언서 동반여행 상품 기획자들로 하여금 고객의 만족도를 이끌어 내기 위한 구체적인 지침을 제공할 수 있을 것으로 기대한다.

주제어

인플루언서 동반여행, 여행상품 선택속성, 인플루언서 특성, 만족도, 추천의도, 브랜드 이미지, S-O-R 프레임워크

B1.4 이차원 고객충성도 세그먼트 기반의 하이브리드 고객이탈예측 방법론

김형수
한성대학교 IT공과대학
hskim@hansung.ac.kr

홍승우
한성대학교 산업경영공학과
seung6858@gmail.com

Abstract - CRM의 하위 연구 분야인 고객이탈예측은 최근 비즈니스 머신러닝 기술이 발전으로 인해 빅 데이터 기반의 마케팅 전략 주제로 더욱 그 중요도가 높아지고 있다. 그러나, 기존 관련 연구는 예측 모형 자체의 성능을 개선시키는 것이 주요 목적이었고, 전체적인 고객이탈예측 프로세스를 개선하고자 하는 연구는 상대적으로 부족했다. 본 연구는 성공적인 고객이탈관리라 모형 자체의 성능보다는 전체 프로세스의 개선을 통해 더 잘 이루어질 수 있다는 가정 하에, 이차원 고객충성도 세그먼트 기반의 고객이탈예측 프로세스 (BiLose-CPP: Bidirectional Loyalty segmentation based Churn Predictive Process)를 제안한다. BiLose-CPP는 양방향 즉 양적 및 질적 로열티 기반의 고객세분화를 시행하고, 고객세그먼트들을 이탈패턴에 따라 2차 그룹핑을 실시한 뒤, 이탈패턴 그룹별로 이질적인 이탈예측 모형을 독립적으로 적용하는 일련의 이탈예측 프로세스이다. 제안한 이탈예측 프로세스의 상대적 우수성을 평가하기 위해 일반적인 범용이탈예측 프로세스들과의 성능 비교를 수행하였다. 글로벌 NGO 단체인 W사의 협력으로 후원자 데이터를 활용한 분석과 검증을 수행했으며, 제안한 BiLose-CPP의 성능이 다른 이탈예측 방법론에 비해 더 우수한 성능을 보이는 것으로 나타났다. 이러한 이탈예측 프로세스는 이탈예측에도 효과적일 뿐만 아니라, 다양한 고객 통찰력을 확보하고, 관련된 다른 퍼포먼스 마케팅 활동을 수행할 수 있는 전략적 프레임워크가 될 수 있다는 점에서 연구의 의의를 찾을 수 있다.

Key Terms - CRM, 고객관계관리, 고객이탈 예측, 고객충성도 세그먼트, 빅 데이터

Acknowledgement: 본 연구는 한성대학교 교내 학술연구비 지원에 의해 수행되었습니다.

1. 서론

성숙기에 접어든 대부분의 업종들이 치열한 경쟁환경에 노출되면서 고객생애가치의 중요성을 인식하고, 이에 따라 신규고객의 확보보다 기존고객의 이탈방지가 더욱 중요한 비즈니스 이슈임을 인지하고 있다 (Hung et al., 2008). 이는 기존 고객을 유지하는 것이 새로운 고객을 확보하는 것에 비해 경제적인 측면에서 훨씬 유리하기 때문인데 (Nie et al., 2009), 실제로 신규 고객 획득 비용은 기존 고객의 유지비용 대비 5~6배에 달하는 것으로 알려져 있다 (Athanasopoulos, 2000). 뿐만 아니라, 고객 이탈은 안정적인 매출의 급격한 감소뿐만 아니라, 밀빠진 독에 물붓기와 같이 일정 수준의 신규 고객을 지속적으로 유치해야 하는 비효율적 경영구조에 빠지게 된다. 이와는 반대로 고객이탈을

효과적으로 방지하여 고객 유지율을 개선시키는 기업은 회사의 수익성 증대뿐만 아니라, 고객만족을 통한 브랜드 이미지 개선에도 긍정적인 효과를 가져오는 것으로 알려져 있다. 즉, 현재와 같은 무한경쟁 환경 속에서는 고객의 이탈을 효과적으로 예측하고, 방지하는 것이 재무적, 비재무적 관점에서 기업에게 효과적이다.

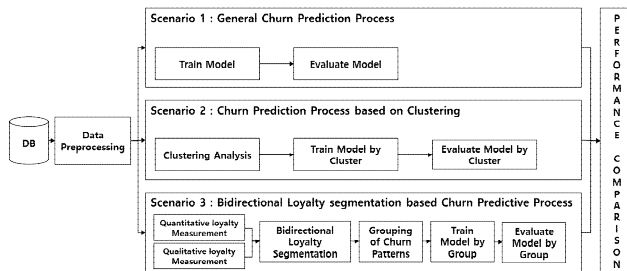
이에 따라, 경쟁이 치열하고, 이탈관리가 시급한 이동통신산업 (e.g., S.A. Qureshi et al., 2013), 금융산업 (e.g., Y. Xie et al., 2008), 게임산업 (e.g., J. Kawale et al., 2009) 등 다양한 업종에서 고객유지율 개선을 목표로 하는 고객이탈예측 연구가 활발히 진행되어 왔다. 기존의 고객이탈예측 연구는 주로 머신러닝과 인공지능 기술을 활용한 예측모형 개발에 초점이 맞추어져 왔으며, 인공지능경망(e.g., S.V. Nath et al., 2003), 의사결정나무(e.g., Lim and Hur, 2006), 로지스틱 회귀분석(e.g., G. Kraljevic et al., 2010), SVM(e.g., Seo 2006), 앙상블 모형(e.g., Lee and Lee 2006; Lee and Kwon 2013) 등 머신러닝 기법들이 사용되어 왔다.

그러나, 이러한 이탈예측 연구는 단순히 다양한 모형의 성능을 비교하거나 (e.g., Kim et al 2002), 이탈예측에 효과적인 변수(feature)를 탐색하거나 (e.g. De Bock and den Poel, 2012), 새로운 앙상블 기법을 개발하는 것 (e.g., Lee and Lee 2006; Lee and Kwon 2013)과 같이 이탈예측 모형 자체의 성능을 개선하는 것에 집중되어왔고, 대부분 전체 고객을 하나의 그룹으로 간주하고 예측모형을 개발했기 때문에 실질적인 활용도 관점에서 한계를 가진다. 왜냐하면 실제 비즈니스상의 고객들은 이질적인 거패턴으로 인해 고객의 행태특성이 다르고, 이에 따른 이탈율도 다르기 때문에 고객 전체를 하나의 고객집단으로 전제하는 것은 무리가 따르기 때문이다(e.g., Kim and Shin, 2012). 이처럼 전체 고객을 대상으로 단일 이탈예측 모형이 작동할 경우, 특정 고객 세그먼트의 고객행태 특성이 달라질 때 변화된 고객특성이 이탈예측 모형에 반영되기 어렵기 때문에 효과적인 고객이탈예측은 불가능해진다.

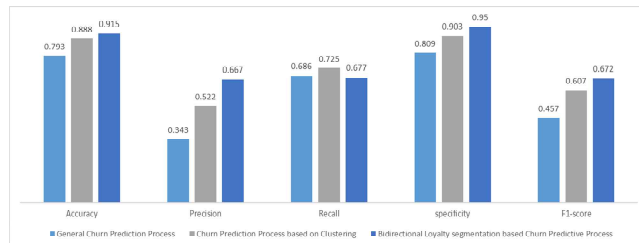
우리는 이러한 관점에서 지속적으로 효과적인 고객이탈 예측을 수행하기 위해서 충성도 관점의 의미 있는 고객 세그먼트를 도출하고, 이를 기반으로 적합한 이탈예측 모형이 개별적으로 운영되어야 한다고 믿는다. 물론 일부 연구에서는 클러스터링 기법으로 고객을 세분화하여 개별 고객그룹에 대한 이탈예측 모형을 적용한 연구는 존재한다 (e.g., Sejoon Oh et al., 2018). 이러한 이탈예측 프로세스는 전체 고객집단을 대상으로 한 단일 이탈예측 모형에 비해 우수한 결과를 가져올 수 있지만, 클러스터링 기법이 투입변수들간의 거리를 기반으로 한 탐색적

그룹핑 기법이기에 때문에 고객의 로열티 증대라는 기업의 전략적 의도가 제대로 반영되지 않을 수 있다는 점에서 여전히 개선의 여지가 남아있다. 이에 본 연구에서는 이차원 고객충성도 세그먼트 기반의 고객이탈예측 프로세스 (BiLose-CPP: Bidirectional Loyalty segmentation based Churn Predictive Process)를 제안하고자 한다. 이러한 연구모형은 비즈니스 머신러닝의 성과가 알고리즘이나 모형 자체의 우수성보다는 전체 머신러닝 프로세스의 우수성에 의해 더 큰 영향을 받는다는 믿음을 기반으로 한다. 본 연구에서는 BiLose-CPP를 제안하고, 기존의 두 가지 범용적인 이탈예측 방법론과 비교하여 모형의 타당성을 입증하고자 한다.

II. Figures and Tables



<그림 1> Customer Churn Prediction Process by Scenario



<그림 2> Process Performance Comparison

<표 1> Process Performance Comparison

III. 참고문헌

Athanassopoulos, Antreas D. "Customer satisfaction cues to support market segmentation and explain switching behavior." *Journal of business research*, Vol.47, No.3 (2000), 191-207.

Chang, M.S., and H. J. Kim. "A Customer Segmentation Scheme Base on Big Data in a Bank." *Journal of Digital Contents Society*, Vol.19, No.1 (2018), 85-91.

Chang, N.S. "Improving the effectiveness of customer classification models: a pre-segmentation approach." *Information Systems Review*, Vol.7, No.2 (2005), 23-40.

Cho, Y. H. and J. H. Bang. "Applying centrality analysis to solve the cold-start and sparsity problems in collaborative filtering." *Journal of Intelligence and Information Systems*, Vol.17, No.3 (2011), 99-114.

Cho, Y. S., M. H. Huh and K. H. Ryu,

"Implementation of Personalized Recommendation System using RFM method in Mobile Internet Environment." *Journal of The Korea Society of Computer and Information* Vol.13, No.2 (2008), 41-50.

De Bock, K. W., and D. Van den Poel, "Reconciling performance and interpretability in customer churn prediction using ensemble learning based on generalized additive models." *Expert Systems with Applications*, Vol.39, No.8 (2012), 6816-6826.

Han, S. L. and A. Maeng, "Effects of Switching Barrier of Mobile Telecommunication Service on Customer Retention And Churn-out" *Journal of Product Research*, Vol 23 (2005), 89-104.

Hung, C., and Tsai, C. F., "Market segmentation based on hierarchical self-organizing map for markets of multimedia on demand." *Expert systems with applications*, Vol.34, No.1 (2008), 780-787.

Hung, S. Y., Yen, D. C., and Wang, H. Y., "Applying data mining to telecom churn management." *Expert Systems with Applications*, Vol.31, No.3 (2006), 515-524.

Joe, D. Y and K. Nam, "SKU recommender system for retail stores that carry identical brands using collaborative filtering and hybrid filtering" *The Journal of Intelligence and Information Systems*, Vol.23, No.4 (2017), 77-110.

Kawale, J., Pal, A., and Srivastava, J. "Churn prediction in MMORPGs: A social influence based approach." 2009 International Conference on Computational Science and Engineering, Vol.4. IEEE, 2009.

Kim, C. N., N. S. Chang, and J. W. Kim, "A Study on the Analysis of Comparison of Churn Prediction Models in Mobile Telecommunication Services" *Asia Pacific Journal of Information Systems*, Vol.12, No.1 (2002), 139-158.

Kim, D. H and K.K Ahn, "A Study on the

Scenario	Accuracy	Precision	Recall	Specificity	F1-score
1. General Churn Prediction Process	0.793	0.343	0.686	0.809	0.457
2. Churn Prediction Process based on Clustering	0.880	0.522	0.725	0.903	0.607
3. Bidirectional Loyalty segmentation based Churn Predictive Process	0.915	0.667	0.677	0.950	0.672

Customer Segmentation Using Machine Learning" *The Society of Convergence Knowledge Transactions*, Vol.6, No.2 (2018), 115-120.

Kim, H. S. and D. Y. Shin, "Developing a customer defection model based on relative defection criteria for non-contractual businesses" *Korean Journal of Marketing*, Vol.27, No.3 (2012), 117-144.

Kim, H. S., Y. L. Bak, and J. M. Lee, "A Personalized Recommendation System Using Machine Learning for Performing Arts Genre" *Information Systems Review*, Vol.21, No.4 (2019), 31-45.

Kraljević, G., and Gotovac, S., "Modeling data mining applications for prediction of prepaid churn

in telecommunication services." *Automatika*, Vol.51, No.3 (2010), 275-283.

Lee, J. S. and J. G. Kwon, "A Hybrid SVM Classifier for Imbalanced Data Sets" *The Journal of Intelligence and Information Systems* Vol.19, No.2 (2013), 125-140.

Lee, K. C., S. J. Kwon, and K. S. Shin, "Analysis of Defection Customer Using Customer Segmentation on Bank" *The Journal of Intelligence and Information Systems*, Vol.7, No.1 (2001), 177-196.

Lim, S. H. and Y. Hur, "Customer Churning Forecasting and Strategic Implication in Online Auto Insurance using Decision Tree Algorithms", *Information Systems Review*, Vol.8, No.3 (2006), 125-134.

Nath, Shyam V., and Ravi S. Behara. "Customer churn analysis in the wireless industry: A data mining approach." *Proceedings-annual meeting of the decision sciences institute*, Vol. 561. 2003.

Nie, G, Wang, G., Zhang, P., Tian, Y., and Shi, Y., "Finding the hidden pattern of credit card holder's churn: A case of china." *International Conference on Computational Science*. Springer, Berlin, Heidelberg, 2009.

Oh, S. J., E. J. Lee, J. Y. Woo, and H. K. Kim, "Constructing and Evaluating a Churn Prediction Model using Classification of User Types in MMORPG." *KIISE Transactions on Computing Practices*, Vol.24, No.5 (2018), 220-226.

Park, S. H. and H. S. Kim, "Design of a Diversity-Based Recommender System for Providing Anti-Churning Rules", *Korea Intelligent Information System Society*, Vol.2011, No.5 (2011), 253-259.

Pham, M. C., Cao, Y., Klamma, R., and Jarke, M., "A clustering approach for collaborative filtering recommendation using social network analysis." *J. UCS*, Vo.17, No.4 (2011), 583-604.

Seo, K. K., "A Study on Customer Segmentation Prediction Model using Support Vector Machine" *Journal of the Korea management engineers society*, Vol.11, No.2 (2006), 97-108.

Shaaban, E., Helmy, Y., Khedr, A., and Nasr, M., "A proposed churn prediction model." *International Journal of Engineering Research and Applications*, Vol.2 No.4 (2012), 693-697.

Tsai, C. F., and Lu, Y. H. "Customer churn prediction by hybrid neural networks." *Expert Systems with Applications*, Vol.36, No.10 (2009), 12547-12553.

Wen, J., and Zhou, W. "An improved item-based collaborative filtering algorithm based on clustering method." *Journal of Computational Information Systems*, Vol.8, No.2 (2012), 571-578.

Xie, Y., and Li, X., "Churn prediction with linear discriminant boosting algorithm." *2008 International Conference on Machine Learning and Cybernetics*. Vol. 1. IEEE, 2008.

Xie, Y., Li, X., Ngai, E. W. T., and Ying, W., "Customer churn prediction using improved balanced random forests." *Expert Systems with Applications*, Vol.36, No.3 (2009): 5445-5449.

B1.5 Users' acceptance toward human-robot interactive service robots : based on YouTube reviews

여방
순천향대학교 경영대학
luvfang6@sch.ac.kr

장수평
순천향대학교 경영대학
zhangxiuping@sch.ac.kr

최재원
순천향대학교 경영대학
jaewonchoi@sch.ac.kr

Abstract – With the advancement of artificial intelligence technology, the human-computer interactive service robot industry has also developed rapidly. Human-computer interactive service robots are robots that have human appearance and action capabilities and can accomplish human service work. Social robots enter our lives as companions and assistants for the elderly and for children with special needs. Robots are used for entertainment, education. Provide more than service automation but act as sales staff, product advisors, shopping assistants for the general public. In recent years, although the research on human-computer interactive service robots has made great progress, there are still many problems. Although the service robot closely resembles a human in appearance but then does not behave like one, there is the danger of the human-robot interaction breaking down. At the same time, infringement of privacy and mechanical malfunction lead to accidents in a harmful way. In this digital age, one of the most powerful promotional tools is the Internet. The Internet has promoted people to express themselves on social media such as Facebook, YouTube, Twitter and Instagram, and everyone can express themselves freely. Analyzing online comment reviews is a common practice as it gives insights on how consumers perceive a product. Therefore, the research aims to uncover user acceptance of human-robot interactive service robots based on YouTube online reviews. By crawling the relevant comments of the YouTube service robot, using word2vec, sentiment classification, and LDA analysis methods to determine the user's acceptance of the human-robot interaction service robot. This research is of great significance to promote users to have a better experience and get more convenience through the understanding of the human-robot interaction service robots so that users can accept the human-robot interaction service robots.

Key Terms – human-computer interactive service robot, LDA, Sentiment classification, word2vec, YouTube

C1 [학술논문] 딥러닝 활용

C1.1 RNN을 이용한 스마트 홈 환경의 난방 제어에 대한 연구

김부건
서울과학기술대학교
IT정책전문대학원
bugeonkim@cvnet.co.kr

김경옥
서울과학기술대학교
산업공학과
kyoungok.kim@seoultech.ac.kr

Keywords : Machine Learning, RNN, Thermostat

I. 서론

스마트 홈 시스템을 이용한 난방 기기 제어 시스템의 활용 빈도가 증가하는 추세이다. 본 논문에서는 기계학습 기반의 온도 조절기 시스템을 제안하였다. 기존 연구에서는 실내 온도 혹은 실외 기상 정보만을 각각 이용하였지만 본 논문에서는 이를 종합하여 이용하였고 또한 기존 논문의 실험 데이터는 시뮬레이터를 이용한 가상의 데이터였지만 본 논문의 실험 데이터는 실제 사용자가 이용한 난방 기기의 실측 데이터이다. 또한, 기존의 연구에서는 난방 기기 예측 기계학습 알고리즘으로 ANN을 활용하였지만 본 연구에서는 난방 기기 데이터의 시계열 특성을 감안하여 RNN 알고리즘을 이용하였다. 이를 통해 기계학습을 활용한 실내 온도 조절기의 예측 가능성을 확인하였고 실제 데이터를 통해 성능을 검증하였다. 실험 결과 제안된 시스템은 사용자의 생활 패턴에 따라 성능의 편차가 존재하지만 일정한 생활 패턴인 경우 우수한 결과를 보였다.

II. 참고문헌

여나영, 마평수, " 유사 날씨 정보를 이용한 건물의 난방 부하 예측 방법," 한국정보과학회 논문지, Vol.25, No.8(2019), 369-376.

안준영, et al., " 동계 전력수요예측을 위한 신경망 모델에 관한 연구," 한국정보기술학회 논문지, Vol.15, No.9(2017), 1-9.

김상훈, et al., " 스마트 홈 환경을 위한 기계학습 기반의 실내 온도 제어," 한국통신학회 학술대회논문집, 2017, 1294-1295.

길소현, et al., " 딥러닝을 이용한 에너지 수요 예측 방법에 관한 연구," 한국통신학회 학술대회논문집, 2016, 1014-1015.

백용규, et al., " 난방시스템 최적 셋백온도 적용시점 예측을 위한 인공신경망모델 개발," 한국생태환경건축학회 논문집, Vol.17, No.3(2016), 89-94.

신동하, et al., " 하계 전력수요 예측을 위한 딥러닝 입력 패턴에 관한 연구," 한국정보기술학회 논문지, Vol14, No.11(2016), 127-134.

김종화, et al., " 순환신경망 모형을 활용한 시계열 비교예측," 한국자료분석학회 논문지, Vol.21, No.4(2019), 1771-1779.

신동하, et al., " RNN과 LSTM을 이용한 주가 예측을 향상을 위한 딥러닝 모델," 한국정보기술학회 논문지, Vol.15, No.10(2017), 9-16.

이창용, 김진호, " 엘만 순환 신경망을 사용한 전력 에너지 시계열의 예측 및 분석," 한국산업경영시스템학회 논문지, Vol.41, No.1(2018), 84-93.

C1.2 전문성 이식을 통한 딥러닝 기반 전문 이미지 해석 방법론

Deep Learning-based Professional Image Interpretation Using Expertise Transplant

김태진 (Taejin Kim)

소속(국문): 국민대학교 비즈니스IT 전문대학원

소속(영문): Graduate School of Business IT, Kookmin University

E-mail: jeung722@kookmin.ac.kr

김남규 (Namgyu Kim)

소속(국문): 국민대학교 경영정보학부

소속(영문): School of MIS, Kookmin University

E-mail: ngkim@kookmin.ac.kr

Corresponding author: Namgyu Kim

77 Jeongneung-ro, Seongbuk-gu, Seoul 136-702, Korea

Tel: +82-2-910-5425, Fax: +82-2-910-4017, E-mail: ngkim@kookmin.ac.kr

국문초록

최근 텍스트와 이미지 딥러닝 기술의 괄목할만한 발전에 힘입어, 두 분야의 접점에 해당하는 이미지 캡셔닝 기술 및 그 활용에 대한 관심이 급증하고 있다. 이미지 캡셔닝은 주어진 이미지에 대해 관련 캡션을 자동으로 생성하는 기술로, 이미지 이해와 텍스트 생성을 동시에 다룬다는 특징을 갖는다. 이미지 캡셔닝은 이미지와 텍스트 데이터를 모두 처리해야 한다는 높은 진입 장벽에도 불구하고, 다양한 활용 가능성으로 인해 인공지능 분야의 연구에서 핵심 분야 중 하나로 자리매김하고 있다. 이에 이미지 캡셔닝의 성능을 다양한 측면에서 향상시키고자 하는 시도가 꾸준히 이루어지고 있으며, 최근에는 이미지를 단순히 정확하게 기술하는 것을 넘어 이미지에 내재된 정보를 더욱 세련되게 전달할 수 있는 고급 캡션을 생성하기 위한 연구도 이루어지고 있다. 이처럼 이미지 캡셔닝의 성능을 고도화하기 위한 최근의 많은 노력에도 불구하고, 이미지를 일반인이 아닌 분야별 전문가의 시각에서 해석하기 위한 연구는 찾아보기 어렵다. 동일한 이미지에 대해서도 이미지를 접한 사람의 전문 분야에 따라 관심을 갖고 주목하는 부분이 상이할 뿐 아니라, 전문성의 수준에 따라 이를 해석하고 표현하는 방식도 다르다. 이에 본 연구에서는 전문가의 전문성을 활용하여 이미지에 대해 해당 분야에 특화된 캡션을 생성하기 위한 방안을 제안한다. 구체적으로 제안 방법론은 방대한 양의 일반 데이터에 대해 사전 학습을 수행한 후, 소량의 전문 데이터에 대해 미세 조정을 진행하는 전이 학습을 통해 해당 분야의 전문성을 이식한다. 또한 본 연구에서는 이 과정에서 발생하게 되는 관찰간 간섭(Interference) 문제를 해결하기 위해 ‘특성 독립 전이 학습’ 방안을 제안한다. 제안 방법론의 실현 가능성을 파악하기 위해 MSCOCO의 이미지 12만 건과 일반 캡션 약 60만 건의 데이터에 대해 사전 학습을 수행하고, 미술 치료사의 자문을 토대로 생성한 ‘이미지-전문 캡션’ 쌍 약 300건의 데이터를 활용하여 전문성을 이식하는 실험을 수행하였다. 실험 결과 일반 데이터에 대한 학습을 통해 생성된 캡션은 전문적 해석과 무관한 내용을 다수 포함하는 것과 달리, 제안 방법론에 따라 생성된 캡션은 이식된 전문성 관점에서의 캡션을 생성함을 확인하였다.

주제어

딥러닝, 전문성 이식, 전이 학습, 이미지 캡셔닝, 인공지능

C1.3 사전 활성화 어텐션 모듈 기반 유방암 이미지 분류 방안

Breast Cancer Classification based on Pre-activated Attention Module

양석우 (Seok Woo Yang)

소속(국문): 가톨릭대학교 심리학전공

소속(영문): Department of Psychology, The Catholic University of Korea

E-mail: josh.yang950@gmail.com

이홍주 (Hong Joo Lee)

소속(국문): 가톨릭대학교 경영학전공

소속(영문): Department of Business Administration, The Catholic University of Korea

E-mail: hongjoo@catholic.ac.kr

Corresponding author: Hong Joo Lee

85 Jibong-ro, Bucheon, Gyeonggi 14662, Korea

Tel: +82-10-3887-1383, Fax: +82-2-2164-4280, E-mail: hongjoo@catholic.ac.kr

국문초록

유방 조영술은 유방암을 진단하고 선별하기 위한 검사 과정으로 유방암의 조기 발견을 목표로 있다. 하지만 유방암 조직 모양의 다양성과 노이즈 등의 문제로 많은 오류가 발생하고 있다. 이를 보완하기 위한 방안으로 전문가의 의사 결정 과정에서 전문의를 보조하는 역할의 CAD(Computer-aided Detection and Diagnosis) 시스템이 제안되었고, 딥러닝의 발전은 기존의 머신러닝 알고리즘의 한계를 극복하며 이미지 판독 오차를 더욱 줄이는데 기여하였다.

최근 딥러닝 기반의 CAD 시스템이 전문의와 비슷한 수행을 보이며, 다양한 알고리즘의 CAD 시스템이 제안되었다. 본 연구는 어텐션 블록을 적용한 CNN 모델의 어텐션 특징지도를 추가 정제함으로써 영상의학 판독 정확도를 개선하는 방법론을 제안한다. 특징지도 정제를 위해 어텐션 블록 뒤에 위치하는 1x1 합성곱 필터 사이즈를 사용한 사전 활성화 구조를 활용하였다.

제안하는 방안을 수행함에 있어 데이터는 Kaggle 오픈 소스 데이터인 DDSM(Digital Database for Screening Mammography)을 사용하였다. 데이터의 클래스는 발견된 병변에 대하여 종괴와 석회화의 진행 정도에 따라 양성과 악성, 총 4 가지의 경우로 분류하였다. 데이터 학습을 위한 딥러닝 모델은 Deep, Residual, Wide, Densely Connected 계열의 CNN 모델 중 8 개를 사용하였다. 데이터의 10 겹 교차 검증 결과, 본 연구의 방법론을 적용한 대부분의 CNN 모델이 그렇지 않은 CNN 모델에 비하여 AUC 점수 지표 측면에서 개선된 성과를 보였다.

주제어

어텐션, 이미지 분류, 유방 조영술, 사전 활성화 구조

C1.4 Algorithms vs. Data? LSTM model based Demand Forecasting Using External Weather

김기태
KAIST 경영대학
kitae.kim@kaist.ac.kr

박성혁
KAIST 경영대학
sunghyuk.park@kaist.ac.kr

Abstract - 수요 예측은 고객 수요를 예측하여 기업 운영상의 리스크 줄이기 위해서 널리 사용되는 방법론이다. 수요 예측 모형의 성능 개선을 위해서는 알고리즘을 고도화 하거나, 모형의 설명력 향상에 도움이 되는 데이터를 추가하는 것이 가능하다. 전통적 수요예측 모형인 ARIMA는 시계열로 표현되는 수요량 데이터 외에, 날씨와 같은 외부 데이터를 대량으로 학습하기 어렵다는 것이 한계점으로 알려져 있다. 최근에 소개된 딥러닝 기반 LSTM 모형은 과거 시점의 수요량 뿐만 아니라 날씨와 같은 외부 변수를 학습하는 데 유리하며, 예측 정확도 역시 높일 수 있는 것으로 알려져 있기 때문에 차세대 수요 예측 방법론으로 주목 받고 있다. 본 연구에서는 수요 예측을 수행함에 있어 알고리즘 고도화에 의한 개선 효과 및 외부 데이터 추가에 의한 개선효과를 구분하여 각 요인이 얼마나 성과를 향상시키는지에 대해 비교하였다. 분석 결과, LSTM 모형을 기반으로 세 가지 날씨 변수를 모두 사용하였을 때, RMSE로 측정해준 모형의 예측력 개선도가 33.4%로 가장 높게 일어났다. 세부적으로, 알고리즘에 의한 개선도는 29.42%, 외부 데이터에 의한 성능 개선도는 6.18%로 일어났다. 이를 통해 알고리즘에 의한 성과와 추가 데이터에 의한 성과를 구분해줌으로써 각 요인 별로 소요되는 비용 대비 ROI 평가가 가능하므로, 딥러닝 모형 개발에 있어 적정 예산 및 기대효과 분석 목적으로 활용 가능하다.

Key Terms - Demand Forecasting, Neural Network, LSTM, External Data, Predictive Analytics

I. 서론

수요 예측은 기업이 미래 시점의 비즈니스를 수행하는데 필요한 자원을 사전에 준비하도록 해주기 때문에 효율성을 높이기 위하여 필수적인 분석활동이다. Deep learning 기반의 LSTM(long short-term memory)모형은 ARIMA(autoregressive integrated moving average)모형과 같은 전통적인 시계열 데이터 분석 방법론보다 더 뛰어난 예측 정확도를 보여주는 것으로 알려져 있으며, 대량의 외부 변수를 학습에 사용하는데 제약이 적어 수요예측 분야에서 재조명 되고 있다 (Siarni-Namini 외, 2018). 오프라인 상에서 비즈니스를 수행하는 기업의 경우, 수요 예측을 진행함에 있어 날씨에 대한 정보를 함께 학습시킬 때, 과거 시점의 시계열 데이터만을 사용하는 전통적인 방식 대비 정확도를 개선하는데 도움이 된다고 알려져 있다 (Ke 외, 2017). Bakker 외(2014)의 연구에서는 모형 개발에 있어서 외부 데이터를 수집하게 되면 별도의 비용이 발생하기 때문에, 데이터 도입 의사결정을 하기 전 단계에서 데이터에 의해 예측 모형의 정확도가 얼마나 개선될 것인지 분석해보는 리스크 평가의 필요성을 강조한다. 그럼에도 불구하고, 실제로

LSTM모형을 사용하여 수요 예측을 수행할 때, 날씨 변수를 사용하는 것이 정량적으로 얼마나 도움이 될 것인지에 대해 명확하게 설명하는 연구는 진행된 바가 거의 없다. 본 연구에서는 수요 예측 과정에서 최신 알고리즘에 의해 예측 정확도가 개선되는 부분과, 외부 데이터 추가에 의해 개선되는 부분을 구분하여 각각이 기여하는 바를 측정하고 비교해보았다. 구체적으로, 수요 예측에 사용되는 데이터의 양을 다르게 설정하고, 알고리즘의 경우에도 수요 예측을 위해 주로 사용되는 ARIMA 모형 및 가장 기초 모형인 MA(moving average) 알고리즘과의 비교를 통해, LSTM모형의 성과를 분석하였다. 이 때, 데이터 관점에서도 외부 날씨 변수를 추가해준 경우와 그렇지 않은 경우로 구분하여 각각 분석해주었다.

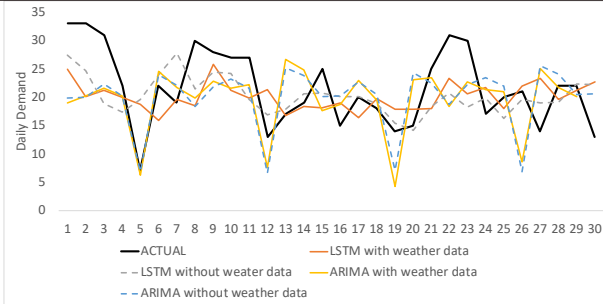
II. 분석 결과 및 결론

본 연구에서 수요 예측을 위하여 기본적으로 사용한 데이터는 과거 시점에 대해 시계열 형식으로 표현되는 수요 데이터에 해당하며, 외부 데이터로는 강수량, 구름량, 풍속에 해당하는 세 가지 날씨 변수를 추가하여 사용 하였다. 수요량의 경우, 날씨가 수요에 큰 영향을 미치는 방문 세차 서비스 신청 이력을 해당 기업에서 제공받아 분석에 사용하였다. 구체적으로, 서울 지역에서 2018년1월1일부터 2020년 3월 5일 동안 수집된 일자 별 방문 세차 예약건수를 수요량으로 사용했다. 모형의 성능은 수요 예측 분야에서 주로 사용되는 RMSE (Root Mean Square Error)을 평가 지표로 선정하고, RMSE가 줄어드는 비율을 개선도로 측정하였다.

분석 결과를 요약한 <표1> 을 보면 LSTM 모형은 가장 단순한 MA 모형 대비 개선도는 29.42%로 확인되었으며, 가장 많이 사용되는 ARIMA 모형과 비교시 개선도가 24.61% (8.901→6.71)로 일어났다. 이는 학습에 사용된 데이터의 양이 많을 수록, 비선형 모형이 예측 정확도를 개선하는데 필요한 정보를 적절히 추출하여 반영한다는 기존 연구(Zhang 외, 1998)를 지지하는 결과에 해당한다. 또한, LSTM모형의 경우 세 가지 날씨 변수를 모두 사용하였을 때 수요에 대한 기본 데이터만 사용해준 경우에 비해 개선도가 6.18% 개선되었다. 이에 비해 ARIMA 모형에서는 이보다 낮은 2.52%의 개선도가 관찰되었다. 이러한 결과는 LSTM 모형이 전통적인 ARIMA 모형에 비해 추가 변수 사용에 대한 제약이 적기 때문에 외부 변수 학습에 유리하다는 기존 연구와(Hochreiter 와 Schmidhuber, 1997; Ke 외, 2017) 동일한 결과를 보인다. <그림1>은 ARIMA 및 LSTM 모형에 대한 추정치를 데이터 별로 보여준다.

<표 1> 알고리즘 별 외부 데이터 사용여부에 따른 성능 개선도 비교

알고리즘	수요량 시계열 데이터 (알고리즘 효과) (데이터 효과)	외부 날씨 변수 추가 (알고리즘 효과) (데이터 효과)
MA	9.507 (0.00%) -	-
ARIMA	8.901 (6.37%) (0.00%)	8.667 (8.84%) (2.52%)
LSTM	6.710 (29.42%) (0.00%)	6.333 (33.39%) (6.18%)



<그림 1> 모형 별 미래 30일간 수요 예측 결과

본 연구는 다음의 시사점을 제시한다. 기업의 데이터 분석가가 현재 사용중인 수요 예측 모형의 성능을 향상시키고자 할 때, 알고리즘을 고도화하거나, 모형 설명력 향상에 도움이 되는 데이터를 추가 확보하려는 시도를 한다. 알고리즘 고도화를 위해서는 LSTM 모델링이 가능한 딥러닝 분석가를 채용하거나, 교육을 받거나, 유료 솔루션을 도입하는 등의 비용이 발생하며, 데이터 추가 확보 과정에서도 데이터 수집을 위한 인력 채용, 서버 확장, 데이터 구독료 등의 비용이 발생한다. 이 때, 각 요인에 의한 성능 개선도를 <표1>과 같이 알 수 있다면, 비용대비 개선 효과를 비교해준 ROI 계산이 가능하다. 향후 연구에서는, 수요 예측 모형의 정확도가 1% 증가할 때마다 운영 비효율 감소에 의해 개선되는 재무적 이익을 계산해주고, 알고리즘 고도화 및 데이터 추가 확보 시나리오 별 ROI 를 진단하는 프레임워크를 개발하고자 한다.

III. 참고문헌

M. Bakker, H. Van Duist, K. Van Schagen, J. Vreeburg, and L. Rietveld, "Improving the performance of water demand forecasting models by using weather input." *Procedia Engineering*. Vol. 70 (2014), 93~102.

Hochreiter, S., & Schmidhuber, J., "Long short-term memory," *Neural computation*, Vol. 9, No.8(1997), 1735~1780.

Ke, J., Zheng, H., Yang, H., & Chen, X. M., "Short-term forecasting of passenger demand under on-demand ride services: A spatio-temporal deep learning approach," *Transportation Research Part C: Emerging Technologies*, Vol. 85, (2017), 591~608.

Siarni-Namini, S., Tavakoli, N., & Namin, A. S., "A comparison of ARIMA and LSTM in forecasting time series," *17th IEEE International Conference on Machine Learning and Applications* (2018), pp. 1394~1401.

Zhang, G., Patuwo, B. E., & Hu, M. Y.,

"Forecasting with artificial neural networks: The state of the art," *International journal of forecasting*, Vol. 14, No.1(1998), 35~62.

C1.5 감정이미지 분류를 위한 CNN과 K-means RGB cluster 앙상블 기법

한기웅
가톨릭대학교

이정현
가톨릭대학교

이홍주
가톨릭대학교

수학과
gksrldnd7897@naver.com

정보통신전자공학부
95743@naver.com

경영학부
hongjoo@catholic.ac.kr

I. 서론

이미지 분류에서 CNN모델을 사용하는 가장 큰 이유는, 이미지의 전체적인 정보에서 각 지역 특징을 추출하여 서로의 관계를 고려할 수 있기 때문이다. 하지만 이미지의 지역 특징이 없는 감정 이미지 데이터는 CNN모델이 적합하지 않을 수 있다. 이러한 감정 이미지 분류의 어려움을 해결하기 위하여 매년 많은 연구자들이 감정 이미지에 적합한 CNN기반의 아키텍처를 제시하고 있다.

본 연구는 사람이 이미지의 감정을 분류하는 기준 중 많은 부분을 차지하는 색감을 이용하여 정확도를 향상시키는 방안을 제안한다. 이미지의 전반적인 색감을 가져오기 위하여, RGB값에 K-means cluster를 적용하였다. 각 이미지 당 많은 부분을 차지하는 대표적인 두 가지 색깔을 추출하여, 각 클래스 별 해당 색깔의 조합이 나올 확률을 가중치 식으로 변형 후, CNN모델의 최종 Layer인 Logsoftmax에 적용하는 앙상블 기법으로 구현되었다. 가중치 식에서 사용한 색감은 색깔 심리학에서 제시하는 16가지 색을 사용하였다.

본 연구에서 제안하는 기법을 수행함에 있어 데이터는 감정 이미지 분류 모델에서 주로 쓰는 ARTphoto, Emotion6, Twitter dataset을 사용하였다. ARTphoto dataset은 총 8가지 감정으로 구분되어 있으며, 긍정적 감정 4가지, 부정적 감정 4가지 데이터로 구성되어 있다. Emotion6 dataset은 총 6가지 감정으로 분류되어 있으며, 각 이미지의 복합적인 감정을 투표 방식을 통해 구분해 두었다. Twitter dataset은 긍정적 및 부정적 감정으로 구분된 이진데이터로 구성되어 있다. 학습에 사용한 CNN모델은 감정 이미지 분류를 실험한 모델 중 가장 좋은 성능을 보였던 Resnet152이며 ILSVRC에서 학습된 pretrained weights를 사용하였다. 성능 평가는 5-fold cross validation으로 CNN모델에 앙상블 적용 전후를 비교하여 검정하였다.

II. 참고문헌

S Liao, J Wang, R Yu, K Sato and Z Cheng, "CNN for situations understanding based on sentiment analysis of twitter data," Elsevier, Procedia computer science, 2017, 376~381.

Understanding the Meaning of Colors in ColorPsychology, 2009. Available at <http://www.empower-yourself-with-color-psychology.com/>. (accessed April, 1. 2014).

K Song, T Yao, Q Ling, and T Mei, "Boosting image sentiment analysis with visual attention," Neurocomputing, 2018.

J Zhang, H Sun, Z Wang and T Ruan, "Another

Dimension: Towards Multi-subnet Neural Network for Image Sentiment Analysis," IEEE, 2019.

K He, X Zhang, S Ren and J Sun, "Deep Residual Learning for Image Recognition," The IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2016, 770~778.

V Gajarla, A Gupta, "Emotion Detection and Sentiment Analysis of Images," Georgia Institute of Technology, 2015.

D1 [학술논문] 지능형 기술응용 I

D1.1 유전 알고리즘(GA)을 이용한 회계부정 탐지모형(SVM) 변수 최적화에 대한 연구

조수현
이화여자대학교
빅데이터분석학 협동과정
soohyuncho7117@gmail.co.kr

신경식
이화여자대학교
경영대학
ksshin@ewha.ac.kr

Abstract - 기업의 투명하고 정확한 공시정보 제공은 자본시장의 건전성을 유지하고 다양한 투자 주체들에게 합리적 투자 의사결정을 지원하기 위해 필수적이다. 본 연구에서는 기업의 사업보고서에서 추출한 기업의 데이터를 이용하여 회계부정 (감리지적기업)을 탐지하는 모형을 구축하고 모형 최적의 변수 집합을 구하기 위하여 유전 알고리즘(Genetic Algorithm)을 이용하여 모형의 성과를 향상시키고자 한다. 또한 기존의 변수 선택기법(t-test, stepwise selection)과 랜덤서치(random search) 기법과 제안 모형의 성과를 비교할 것이다.

Key Terms - 변수 최적화, 유전 알고리즘, 회계부정 탐지
I. 서론

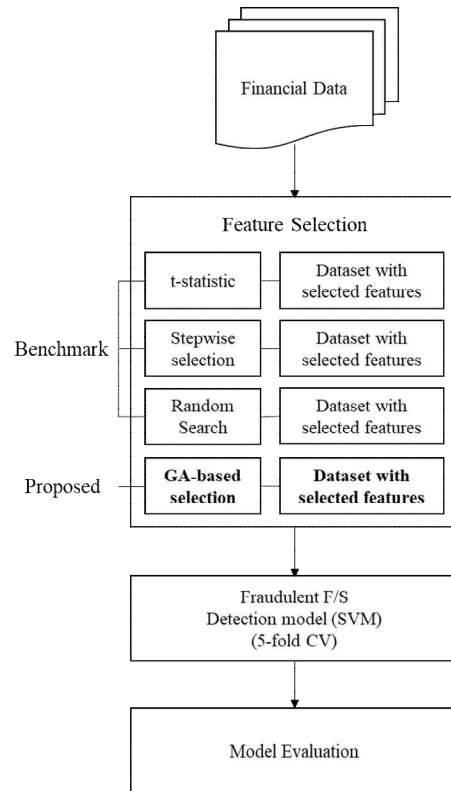
최근 기관 및 개인투자자 등 투자 주체들이 주식투자에 대한 관심이 증가하면서 상장 기업의 투명한 정보공개의 중요성이 강조되고 있다. 기업의 공시정보는 투자자들이 합리적 투자결정을 할 수 있도록 정보를 제공하는 데 목적이 있다. 미국의 엔론(Enron) 사건을 비롯하여 우리나라의 대우 분식회계 사건, 모뉴엘 분식회계 사건 등 기업의 회계부정 사례에서 알 수 있듯 기업의 부정확한 기업 정보제공은 투자자뿐만 아니라 주식시장, 산업 및 경제까지 방대한 영향을 미칠 수 있다. 따라서 불공정하고 부실한 공시를 하는 기업에 대해서는 적발 후 제재를 가하고 있다.

우리나라의 경우 일반적으로 리스크 등을 감안하여 표본을 추출하여 금융감독원에서 재무제표 심사절차를 통하여 회계 위반 사례를 적발한다. 그리고 회계위반사항에 대하여 당국의 수정권고를 따르지 않을 시 감리를 실시하고 증권선물위원회에서 제재 등을 결정하여 위반사항 등을 공시하는 절차를 따른다. 이렇듯 회계감리 절차는 많은 전문가의 노력과 시간이 투입되는 과정이다. 따라서 근래에는 데이터마이닝, 머신러닝 등 분석 기술을 이용하여 회계부정을 탐지하는 모형에 대한 연구가 활발해지는 추세이다. 많은 연구에서는 기업이 공시하는 사업보고서의 정보를 대상으로 다양한 데이터마이닝 기법을 적용하여 회계부정을 적발하는 연구를 수행하였다. (Dutta et al.2017; Ravisankar et al., 2011; Hajek and Henriques ,2017)

대부분의 연구는 데이터마이닝, 머신러닝 기법을 사업보고서에서 추출한 다양한 변수와 데이터에 적용한 연구로 실제 모형구축 시 중요한 변수 선정에 관해서는 큰 관심을 가지지 않았다. 그러나 머신러닝 기법을 이용한 모형 구축에서의 변수선정은 모형의 성능에 영향을 미치는 중요한 단계이다. 본 연구에서는 우리나라 KOSPI, KOSDAQ 상장 기업 중 감리지적 기업과 건전

기업 데이터를 대상으로 유전 알고리즘(Genetic Algorithm)을 이용한 변수 최적화를 통한 회계부정 탐지 모형(SVM)을 구축하고자 한다.

II. Figures and Tables



<그림 1> 유전 알고리즘(GA)을 이용한 회계부정 탐지모형 변수 최적화 연구모형 flowchart

III. 참고문헌

- Dutta, I., Dutta, S., & Raahemi, B. (2017). Detecting financial restatements using data mining techniques. *Expert Systems With Applications*, 90, 374-393.
- Hajek, P., & Henriques, R. (2017). Mining corporate annual reports for intelligent detection of financial statement fraud – A comparative study of machine learning methods. *Knowledge-Based Systems*, 128, 139-152.
- Ravisankar, P., Ravi, V., Raghava Rao, G., and Bose, I. (2011). Detection of financial statement fraud and feature selection using data mining techniques. *Decision Support Systems*, 50(2), 491-500.

D1.2 제로에너지타운 에너지클라우드 데이터 보안 기능 고도화 연구

전민재
고려대학교
컴퓨터융합
소프트웨어
학과
93jmj@kore
a.ac.kr

곽재혁
고려대학교
컴퓨터정보
학과
wlsansdl123
@korea.ac.k
r

이대국
고려대학교
컴퓨터정보
학과
daeguklee@
korea.ac.kr

김민성
고려대학교
컴퓨터정보
학과
cprog95@k
orea.ac.kr

김이강
고려대학교
자연과학연
구소
kimyikang
@korea.ac.k
r

윤석호
고려대학교
컴퓨터융합
소프트웨어
학과
bluepig@ko
rea.ac.kr

조충호*
고려대학교
컴퓨터융합
소프트웨어
학과
chcho@kore
a.ac.kr

요약

본 논문에서는 건물 및 에너지 소비현황, 에너지 생산/운영 설비등의 다양한 중요한 데이터가 존재하는 제로에너지타운 에너지클라우드(Zreo energy Town Energy Cloud: ZTEC)의 보안 기능을 강화하기 위한 방안에 대하여 기술한다. 제안하는 보안 기능을 위해 시스템 구성도 및 보안 시나리오를 제시한다. 또한, 오픈소스 기반의 보안 시스템을 개발하여 ZTEC 시스템에 적용하였다.

Key Terms

제로에너지 타운 에너지클라우드(ZTEC), 가상사설망(VPN), 오픈소스, SoftEther VPN

1. 서론

제로에너지타운 에너지클라우드에는 학교, 어린이 집, 가정집 등의 다양한 건물들에 에너지 소비, 생산, 저장 데이터가 존재하며 이 데이터는 에너지 수요 및 공급 분석 등에 활용된다[1]. 별도의 인증 없이 자유롭게 데이터에 접근이 가능할 경우, 데이터 위변조가 일어날 우려가 있으며 ZTEC 시스템 운영 및 관리에 차질이 생길 수 있다. 사용자 인증을 통해 보안을 강화 해도 다른 사람의 인증을 도용하여 시스템 보안에 위협을 가할 수 있다.

따라서 본 논문에서는 위의 문제점들을 예방하기 위해 발생할 수 있는 보안 위협 사항을 바탕으로 보안 시나리오를 제시하였다. 또한, 시나리오에 맞는 시스템을 설계하고 오픈소스 VPN(Virtual Private Network)인 SoftEther(Software Ethernet) VPN을 기반으로 시스템을 구축하였다.

2. 관련연구

2.1 제로에너지타운 에너지클라우드(ZTEC)

ZTEC은 에너지 소비, 생산, 저장 자원들을 효율적으로 관리하고 건물간 에너지 공유를 위한 클라우드 기반 분산자원 시스템이다. 또한 ZTEC의 목적은 개별 제로에너지빌딩의 높은 건축 비용의 한계를 극복하기 위하여 다수의 빌딩으로 구성된 제로에너지타운을 대상으로 에너지를 공유하는 에너지 클라우드를 구축하고, 에너지 수요공급 최적화를 통하여 제로에너지타운 기반의 경제적인 제로에너지빌딩을 구축, 운영 및 관리하는 것이다.

2.2 SoftEther VPN[2]

일본 츠쿠바 대학에서 만든 멀티 프로토콜 기반 오픈소스 VPN 서버 소프트웨어이며 기존 VPN 서버와 다른 세가지 장점이 있다. 첫째로,

단일 VPN서버에서 여러가지 VPN 프로토콜을 지원한다. 두번째로, 사용자 관리 및 네트워크 기능을 가상화하여 멀티테넌시 가상 호스팅에 필수적인 기능을 제공한다. 마지막으로, 여러 OS와 호환이 가능하여 이식성이 좋다. 이 뿐만 아니라, VPN 기능을 제공하는 장비가 따로 있지만 가격이 비싼 반면에 VPN 서버 도입 비용만으로 VPN 서버를 구축할 수 있는 장점이 있다.

3. 시나리오를 통한 데이터 보안 기능 구현

3.1 제안하는 가상사설망 구성도

ZTEC에 적합한 데이터 및 접근에 대한 사용자 요구사항, 시스템 요구사항 및 그에 따른 보안 시스템을 구현하기 위하여 다음과 같은 구성을 제시한다.

- VPN 서버: 공용 네트워크와 연결이 되어 있으며 유일하게 데이터 서버와 통신할 수 있다.
- 데이터 서버: VPN 서버를 제외한 모든 접근이 차단되어있으며 에너지 소비 및 생산 데이터가 있다.
- SoftEther VPN 클라이언트: 클라이언트 실행 시 로그인 기능을 구현하고 인증된 사용자가 아닌 다른 사람이 로그인에 필요한 정보를 이용하여 VPN에 접속하는 것을 막기 위해 클라이언트 종료 후 로그인에 필요한 정보를 삭제한다.



<그림 1> 가상사설망 구성도

3.2 실험을 위한 보안 시나리오 및 구현

실험 목적은 접근 시나리오를 만들어서 시스템의 불법적인 접근 차단하고 이를 관리할 수 있다는 것을 확인하여 제안하는 구성도를 검증하려 한다. 구성도를 검증하기 위해 시나리오를 3개 만들었다. 시나리오는 다음과 같다.

- 1) SoftEther VPN에 로그인하지 않고 데이터 서버에 접근
- 2) SoftEther VPN에 로그인 실패
- 3) SoftEther VPN에 로그인하고 데이터 서버에

접근
구성한 시나리오를 바탕으로 서버를 구성하고
SoftEther VPN 서버 프로그램의 수정 및 client
프로그램에 편의성을 개선하였다.



<그림 2> 수정된 가상사설망 접속 프로그램

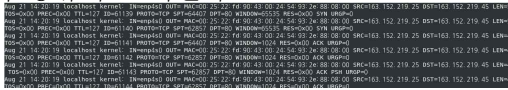
3.3 실험 결과

시나리오 별 결과를 확인해보기 위한 방법으로
표1에 적힌 내용을 수행하였다.

<표 1> 시나리오 체크 리스트

시나리오	VPN로그인시도	VPN로그인	로그
첫번째	X	X	방화벽
두번째	O	X	VPN
세번째	O	O	VPN

1) VPN에 로그인하지 않고 데이터 서버에 접근
이 부분은 VPN 서버 로그가 남지 않기 때문에
데이터 서버에 저장되는 로그를 직접
확인해야한다. 짧은 시간(1초)에 여러 번 호출
될 경우 로그를 저장하도록 한다. 로그에
표시되는 정보는 출발지 IP, 도착지 IP, 접근 시간,
데이터 종류가 표시된다.



<그림 3> 방화벽 접근에 대한 로그

2) VPN에 로그인 실패
VPN에 로그인 실패 혹은 성공 했을 때
VPN서버에 로그인 시도 시간, 로그인 유형,
로그인 시도한 IP, 로그인 시도한 아이디, 로그인
시도한 컴퓨터의 정보 등 다양한 정보가
표시된다.

3) VPN에 로그인 성공
VPN에 로그인 성공했을 때 실패한 경우와
마찬가지로 시간, 로그인한 IP, 로그인한 아이디,
로그인 시도한 컴퓨터의 정보 등 다양한 정보가
표시된다.



<그림 4> VPN 로그

4. 결론

본 논문을 통해 ZTEC의 데이터 보안기능을
향상 시키기 위한 방안을 제시하고, 오픈소스 기반
VPN을 통해 시스템을 구현하여 실험을
진행하였다. 특별한 도구를 사용하지 않고도
오픈소스를 활용해 보안 기능을 강화하고
비정상적인 접근을 제어할 수 있다는 것을
확인하였다. 또한, 오픈소스이기 때문에 소규모
단체 혹은 개인이 운영하는 서버의 보안 강화를
위해서 사용이 가능할것으로 예상된다.

5. 참고문헌

[1] Nam H. S., S. J. Lee, C. H. Shin, J. Y. Kim
and Y. K. Jeon, "Design of Energy Cloud Based
Distribution Resource Management System," The
Journal of Korean Institute of Communications and
Information Sciences, 567-568

[2] Novoridai Yuyu, Shinjio Yasushi, Sato
Satoshi, "SoftEther VPN Server: Multi-Protocol
Cross-Platform open source VPN server, Computer
Software 32.4(2015): 4_3-4_30

D1.3 엔진 데이터 구분 방법론을 통한 실시간 실린더 최대압력 및 IMEP 최적화 ELM 예측 비교

이병석
울산과학기술원 융합경영대학원
eotaor@unist.ac.kr

성노윤
한국과학기술원 경영대학
nyseong@kaist.ac.kr

남기환
한국과학기술원 경영대학
namkh@kaist.ac.kr

Abstract - 이 논문은 H중공업에서 개발한 H엔진 DF 타입 중 스파크 점화 엔진과 흡사한 가스 모드에 대한 실시간 최대 연소압과 IMEP를 예측하기 위해 Extreme Learning Machine(ELM) 모델을 사용한다. 데이터는 발전엔진인 H엔진에 정격 RPM인 1000RPM에서 부하(Load) 20~100%까지 연속된 65000개를 사용하였다. ELM 모델은 SLFN을 기반으로 나온 모델이며, SLFN의 Backpropagation 방법에 기반한 고전 방식의 많은 반복 단계로 인한 장시간의 러닝 시간을 보완하기 위해 채택된 수학적 모델이다. 엔진 연소압 예측에 있어 ELM 모델의 정확도 상승을 위해 엔진 데이터를 전처리 및 데이터 구분(labeling)을 통한 예측 방법론을 제시하고 기존 방법론과 비교하였다. 데이터 구분 방법론은 ELM 모델 결과를 통한 Error 값의 labeling을 거친 후 Decision Tree Classification으로 Data를 그룹화 및 최적화 한 것이다. 제시된 방법론은 기존 방법론을 이용한 예측 결과보다 비약적인 성능 향상 결과를 보이며, 실제 엔진에서 연소압 센서를 대체할 수 있는 잠재력을 보여준다.

Key Terms - P-max prediction, IMEP prediction, PID control, Decision Tree Classification, Gradient boosting regression, Extreme learning machine.

I. 서론

모든 선박엔진 메이커는 해당 기구의 규제를 만족하며, 엔진에 성능 향상을 위한 컨트롤 로직 개발에 주력해야 한다. 환경 규제를 만족하기 위해서는 엔진 가속 성능과 배출 가스에 대한 측정이 필요하며, 측정 시간에 대한 노력을 절감하기 위해 연소 실린더 압력을 통해 예측방법이 연구되고 있다(Henningsson et al., 2012). 현재는 실린더 압력 센서로부터 도시 평균 유효 압력(indicated mean effective pressure, IMEP)과 실린더 최대 연소압 (cylinder peak pressure, Pmax)을 계속하여 질소산화물(NOx) 배출량을 조절하는 기술이 적용되었고, IMO 규제를 만족하기 위한 엔진 컨트롤 시스템개발이 활발하게 이루어지고 있다(Greve, 2014). 하지만 선박 발전용 중형엔진에 경우 실린더 최대 연소압이 200bar 이상 형성되어 계측 센서 파손과 정확도 저하 현상이 빈번하게 발생하고 있으며, 이에 따라 엔진 컨트롤 정확도에 부정적인 영향을 주고 있다. 또한 비용도 일반 압력 센서보다 10배 이상 비싸기 때문에 주변 센서 데이터로부터 신속하고 정확한 실린더 최대 연소압 및 IMEP 예측기법의 도입이 절실하다.

선박용 중형 엔진 메이커는 가스와 디젤연료를 Mode별로 사용할 수 있는 Duel fuel 엔진을 개발했다. 모드 중 가스 연료를 사용하는 경우 Micro Pilot을 이용하여 엔진 점화를 발생시킨다.

이 경우는 스파크 점화 가솔린 엔진과 흡사한 엔진 컨트롤이 적용되며, 이때 연소압은 비선형(nonlinear) 적인 결과물이 나온다. 따라서 가스 모드에 순간 연소압 예측은 일반적인 선형 모델로는 예측이 힘들다. 이에 따라 예측 정확도를 높이고 예측 데이터의 비선형성을 줄이고 최적화 시킬 수 있는 방법을 찾아야 한다.

이에 따라 이 논문에서는 H중공업에서 개발한 H엔진 중 Duel Fuel 엔진에서 연소압 예측이 가장 힘든 Gas Mode 데이터를 이용하여, 연속된 부하구간[20%~100%]내 순간 최대 연소압과 도시 평균 유효압력(IMEP)을 엔진 데이터 그룹핑을 통한 예측 최적화 방법론과 비 그룹핑 방법론을 비교 제시한다.

II. 연구 방법

예측 정확도를 높이고 입력 데이터 차원을 줄이기 위해(Dimension reduction) PCA(Principal component analysis) 방법론을 사용한다(Devasenapati et al., 2010). PCA는 고차원에 데이터를 직교 변환하여 차원수를 줄일 수 있다. PCA는 데이터를 한 개의 축으로 사상 시킬 때 분산이 가장 커지는 축을 첫 번째, 두 번째로 커지는 축을 두 번째 주성분으로 놓이도록 새로운 좌표계로 데이터를 선형 변환한다. 따라서 데이터의 정보를 가지고 서로 간 상관관계(Correlation)가 낮은 새로운 데이터가 생성된다.

이렇게 생성된 데이터와 스피드 컨트롤 변수로 추출된 데이터는 엔진 실시간 부하를 예측하는데 적절한 입력 변수가 되고, 예측에 추가될 로드 데이터를 Gradient Boosting Regression 모델을 이용하여 예측 및 추가했다. Gradient Boosting 방법은 모델 생성 후 예측 오류를 손실함수(Loss function)으로 정량화 후 다음 모델을 보완하는 방식으로 이와 같은 과정을 반복하며 모델을 생성하는 방법론이다.

III. 데이터

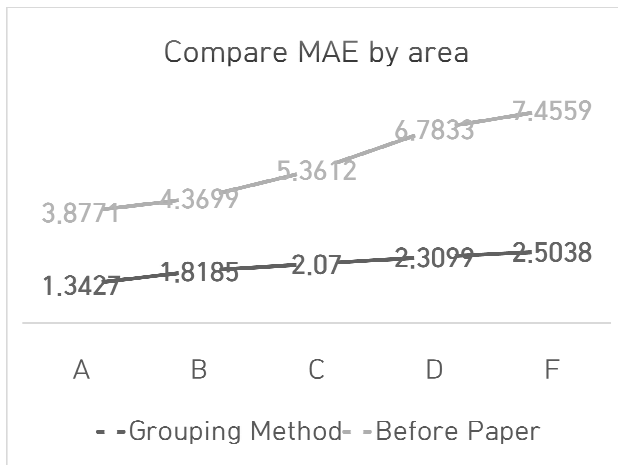
본 연구의 목적을 위해 H엔진 부하를 20~100%까지 연속된 65000개 데이터 세트를 수집했으며, 데이터는 120ms 간격으로 한 개 실린더 사이클 시간과 동일하게 측정했다. 발전기 엔진으로써 RPM은 정격 1000RPM을 유지하고, 이들 구성으로 연소실 내부에 실린더 압력을 측정하였다. 예측에 사용한 추출 데이터는 총 31개이며, 이중 3개를 제외한 데이터가 선급에서 요구하고 있는 Safety 관련한 센서로 구성되어 있다. 센서 종류는 200도 까지 측정 가능한 RTD Sensor, 800도까지 측정가능한 Thermocouple Sensor, 각 Range 별 Calibration 되어있는 Pressure Sensor로 구성되어 있다. 센서의 위치는 크게 Control & Charge Air, Lubrication Oil, Fuel Oil, Exhaust Gas 로 나뉜다. 나머지 3개

입력 항목은 가스모드 운전 시 제어와 관련된 변수(Gas inlet valve timing, duration, Cylinder Duel Valve timing) 이다. 추가로 예측에 사용된 4개 데이터는 정격 1000 RPM과 실제 출력 RPM의 차이 값을 이용해 PID Error Data 3개와 Gradient Boosting Regression모델로 예측한 Load Data 1개 이다. 본 연구에 가장 큰 장점은 실제 양산 엔진에 테스트 데이터를 이용한 데이터이며, 데이터 예측 모델은 실제 엔진 제어기에 삽입 후 제품화로 이어질 수 있다.

IV. 연구 결과

Grouping Method는 본 연구에서 제안된 모델을 의미하고 Before paper는 가장 최근 발표된 연구를 의미한다. Grouping Method와 MAE를 누적 구간별로 비교할 경우 확연한 차이를 <그림 1>처럼 볼 수 있다.

최근 발표된 연구에서는 인기 있는 SLFN 모델에서의 초기 가중치에 대한 종속성, 늦은 수렴시간이라는 단점을 보완한 ELM(Extreme learning machine) 기법을 이용하여 연소압을 예측하였다. 하지만 엔진에 걸리는 로드 및 스피드 변화에 따른 실시간 최대 연소압 및 IMEP를 예측하기에는 단일 ELM 모델로는 한계가 존재한다.



<그림 1> Comparison of method Grouping Method and classic method MAE.

이 논문에서는 Decision tree Classification through ELM Error Feedback labeling Method를 이용한 Grouping 방법을 통해 Grouping Data를 개별적으로 최적화한 예측법을 제시한다. 이 논문에서 제시한 방법론의 결과는 기존 연구들 예측방법으로 찾기 힘들었던 엔진 로드별 연소압 흔들림 추세를 예측 가능했으며, 해당 결과는 NOx 제어를 위한 연소압 데이터로 적용 가능성을 충분히 보였다. 미래 연구에는 실제 예측 연소압을 통해 엔진 성능을 개선시킬 수 있는 AI 컨트롤을 적용해 볼 예정이며, 연소압 예측을 통해 연비 예측, 환경 물질에 대한 예측의 정확도를 높이는 연구 기회가 존재한다.

IV. 참고문헌

Henningsson M, Tunestal P, Johansson R. A Virtual Sensor for Predicting Diesel Engine Emissions from Cylinder Pressure Data 2012 IFAC

Workshop on Engine and Powertrain Control, Simulation and Modeling.

NOx Control With Cylinder Pressure based engine control system. Greve2014_Article.

Devasenapati SB, Sugumaran V, Ramachandran KI. 2010 Misfire identification in a four-stroke four-cylinder petrol engine using decision tree Expert Systems with Applications (37), pp. 2150-2160.

D1.4 국민 참여 방식의 SNS 기자단 데이터베이스시스템 구축방안

Title of the paper in English: A Study on the Establishment of a Database System for SNS Reporters in a Public Participation Method

정기혁 (Ki hyoek Jeong)

소속(국문): 고려대학교 컴퓨터정보학과

소속(영문): School of Computer Information, Korea University

E-mail: knp6464@korea.ac.kr

한울 (wool Han)

소속(국문): 상명대학교 역사콘텐츠학부

소속(영문): School of history, Sangmyung University

E-mail: good930309@hanmail.net

국문초록

최근 여러 정부 기관에서 기자단을 운영하고 있다. 하지만 코로나 19로 모든 행사가 취소되자 기자단 활동에 큰 차질이 생겼다. 기자단은 SNS상에서 다양한 콘텐츠를 제작하지만, 기관의 행사를 취재하여 알리는 역할을 담당했기 때문이다. 행사 자체가 취소되다 보니, 기사 주제가 중복되는 현상이 빈번히 발생하고 있다. 본 연구는 이러한 상황을 해결하기 위해 기자단 활동을 데이터베이스화하는 방안을 제시하였다. 데이터베이스시스템에서 기자단과 담당자가 기사 작성과정을 모아 한 번에 살펴볼 수 있어 주제의 중복을 피하며, 다양한 피드백을 관찰할 수 있을 것이다.

주제어

국민기자단, 데이터베이스시스템, 시민기자단, 공공데이터법,

E1 [학술논문] 지능형 기술응용 II

E1.1 Siamese Attention LSTM Network를 이용한 문장 ↔ 관련법령 비교 시스템 구축 방안 연구

정지인
연세대학교 공학대학
inny1987@yonsei.ac.kr

김민태
연세대학교 공학대학
iammt@yonsei.ac.kr

김우주
연세대학교 공학대학
wkim@yonsei.ac.kr

Abstract - 본 논문에서는 **군내 방산 관련 문서에서 등장할 가능성이 있는 문장과, 이와 관련된 법령 조항과의 유사도를 비교하여 위법 위험 여부를 구분하는 딥러닝 모델을 제안한다.** 모델은 Self-Attention과 Long Short-Term Memory(LSTM)을 결합한 Siamese Attention LSTM 모델이며, 직접 제작한 데이터셋인 **母문장(실제 법령 조항)과 子문장(母문장에서 파생된 변형 문장) 3,420쌍을 대상으로 실험을 수행한 결과 상당히 높은 정확도로 위법 위험 여부를 측정할 수 있었다.**

Key Terms - Siamese network, 법령 비교 시스템, 방위사업, 방위사업법

I. 서론

국방부는 강한 국방력 건설을 위해 방위력개선 사업을 추진중에 있으며, 방위사업관리를 통해 매년 10조원 이상을 집행하고 있다. 2020년의 방위력개선비는 국방예산 50.2조원 중 33.3%인 16.7조원을 차지함으로써 적지 않은 비중을 차지하고 있다. 이러한 방위사업은 국민의 생명과 재산뿐만 아니라 국가의 안보와도 직결되는 만큼 전문가들에 의해 매우 투명하고도 효율적으로 이루어져야 함은 분명하다.

2006년 방위사업청의 개청으로, 기존에 법적 근거 없이 국방획득관리규정에 의해 추진되던 방위사업이 「방위사업법」 제정을 통해 법률에 의해 추진되기 시작했다. 방위사업관리에 적용되는 관련 법 및 규정은 2006년 129개였으나, 2015년에는 174개로 35% 증가되었다. 이러한 상황에서 방위사업관리를 수행하는 담당자는 많은 법과 규정 때문에 업무를 수행하는 데 많은 어려움을 겪고 있다. 업무를 추진한 이후에 모르는 관련 규정이 있다는 사실을 인식하는 경우도 많은 것으로 알려지고 있다.

특히 방위사업 관계 법령(법률, 대통령령, 국방부령) 건수(12건)에 비하여 방사청 소관 행정규칙(훈령, 예규, 고시) 건수가 162건으로 매우 많은 것은 전문적이고 기술적인 특성상 행정규칙을 통해 보다 세부적으로 명시할 필요가 있기 때문이다. 하지만 그에 비례하여 방위사업관리는 더욱 어렵게 되었고 때로는 어떤 법과 규정을 적용해야 하는지도 모르는 상태에서 업무를 담당하는 경우도 있다. (Kim Sun-young, 2017, p. 77)

또한, 법령은 문장 내에서 단어가 하나만 잘못 되더라도 문장의 의미가 크게 바뀌는 아주 예민한 문장으로써 이는 방위사업에서 매우 심각한 문제를 유발할 수 있다. 그럼에도 불구하고 방위사업 추진 간 실무자들은 관련 법령을 직접 손으로 찾아가며 업무를 추진 중이며, 이러한 업무 방식은 많은 시간과 노력이 소모될 뿐만 아니라 법령 임의해석 및 오해석(誤解釋)의 위험이 존재함으로써 치명적인 국고의 손실과

직결될 가능성이 있다.

따라서 본 연구에서는 방위사업 관련 실무자들이 각종 계획문서 및 보고문서 등 주요문서 작성시 실시간으로 참고할 수 있는 자동화된 ‘방위사업 문서 ↔ 관련법령 간 비교 시스템 구축’을 위한 Siamese Attention LSTM 모델을 제안하려고 한다. 이 모델을 통해 방위사업 문서 내에서 등장할 가능성이 높은 문장(子문장)을 이와 관련된 관련법령의 해당 조항(母문장)과 1:1로 비교해 보았으며, 위에서 언급한 바와 같이 두 문장의 전체적인 구조가 매우 유사하더라도 의미의 차이를 유발하는 작은 차이점도 인식하여 관계법령상 위법 위험 여부를 상당히 정확하게 계산할 수 있음을 확인하였다.

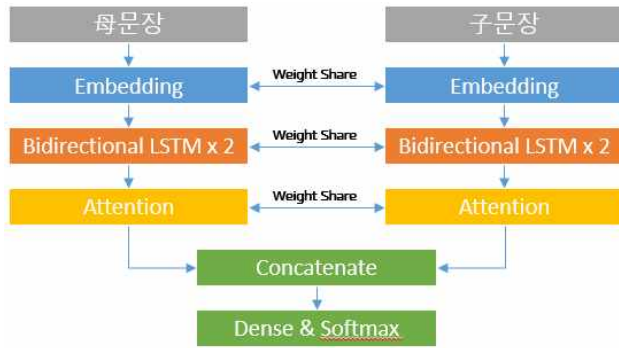
II. Data and Model

본 연구에서 사용한 데이터 셋은 <표1>과 같다. 방위사업 관련법령(방위사업법, 시행령, 시행규칙) 안에서 주요 조항 83개(母문장)를 선정하였으며, 각 1개 조항(母문장) 당 군 문서에서 변형되어 등장할 가능성이 있는 약 30~50개의 유사 문장 (子문장)을 작성하여 실험했다.(총, 3,420 pair)

<표 1> 데이터 셋 세부내용.

구분	母문장	子문장	위험여부
1	국방부장관 및 방위사업청장은 분석·평가 결과 중 총사업비 오천억원 이상의...	국방부장관은 분석·평가 결과 중 총사업비 오천억원 이상의...	1
		국방부장관 및 방위사업청장은 분석·평가 결과 중 총사업비 육천억원 이상의...	0
		합동참모의장은 분석·평가 결과 중 총사업비 오천억원 이상의...	0
.....			
2	국방부장관은 방위력개선사업분야에 관한 중기계획을 대통령의 승인을 얻어 수립한다.	중기계획은 대통령의 승인을 얻어 국방부장관에 의해 수립된다.	1
		중기계획은 국방부장관의 승인을 얻어 방위사업청장에 의해 수립된다.	0
.....			
3	합동참모의장은 방위력개선사업의 소요에 대하여 무기체계 등의 소요를 결정한다.	합동참모의장은 방위력개선사업의 소요를 결정하지 않는다.	0
.....			
4	방위사업청장은 국내에서 생산된 군수품을 우선적으로 구매한다.	방위사업청장은 국내에서 생산된 군수품을 우선 구매한다.	1
.....			
15	전력화평가는 무기체계를 최초로 생산하거나 구매하여 배치한 후 일년 이내 실시하되,	전력화평가는 무기체계를 최초로 생산하거나 구매하여 배치한 후 일년 이후 실시하되,	0
.....			

또한 다음의 <그림 1>은 모델의 구조도이다.



<그림 1> 모델 구조도

III. 참고문헌

Hochreiter, S., and Schmidhuber, J., "Long Short-Term Memory," Neural Computation, 1997.

Kim Sun-young, "Theory and Practice of Defense Acquisition," Bookorea, Inc. Korea, 2017, 77.

Mueller J, Thyagarajan A, "Siamese recurrent architectures for learning sentence similarity," Thirtieth AAAI Conference on Artificial Intelligence, AAAI Press, 2016.

Lin, Zhouhan, et al., "A structured self-attentive sentence embedding," arXiv preprint arXiv:1703.03130, 2017.

Mintae Kim, Yeongtaek Oh, Wooju Kim, "Sentence Similarity Prediction based on Siamese CNN-Bidirectional LSTM with Self-attention," Journal of KIISE 46(3), 2019.

E1.2 계량 측정을 위한 문자 인식 기술 구조

이교혁
연세대학교 기술경영학협동과정
kyohyuk.lee@yonsei.ac.kr

김우주
연세대학교 산업공학
wkim@yonsei.ac.kr

Abstract - 본 연구에서는 계량 자동 측정을 위하여 계량기 이미지를 획득하고 광학 문자 인식 기술을 적용하여 전기, 수도, 가스 등의 계량을 자동으로 수행하는 시스템 구조를 제시하였다. 광학 문자 인식을 위해서 CNN (Convolutional Neural Network)을 이용한 문자 영역 검출 및 CRNN (Convolutional Recurrent Neural Network)을 이용한 문자열 인식 기술을 적용하였으며, 고속 병렬 처리를 위해 GPU 클라우드에 기반하여 시스템을 구축하였다. 전체적인 시스템은 계량기 이미지를 일반 스마트폰 사양의 모바일 기기를 이용하여 획득한 후 LTE 네트워크를 통해 클라우드 서버로 전송하고, 클라우드 서버에서 광학 문자 인식을 수행하는 구조로 이루어져 있다.

Key Terms - 계량 자동 측정, 광학 문자 인식, 합성곱 신경망, 합성곱 순환 신경망

I. 서론

전기, 수도, 가스 등 각종 계량기의 사용량을 측정하기 위해서는 디지털 계량기를 설치하거나, 운영 인력이 직접 확인하여 수작업으로 기입하는 방식을 사용해야 한다. 디지털 계량기의 경우 NB-IoT 등과 같은 machine-to-machine 통신 기술을 적용하며, 통신 모듈 구동을 위한 전력을 배터리로부터 공급 받는다. 배터리는 수명이 정해져 있으므로, 주기적으로 배터리를 교체해야 한다. 배터리가 방전될 경우, 계량 측정이 수행되지 않을 수 있는 문제점이 있다. 수작업 측정의 경우, 매 월 많은 수의 운영 인력이 모든 계량기 위치에 방문하여 수작업으로 입력해야 하므로, 운영 효율 측면에서 매우 비효율적이다. 이러한 문제점을 해결하기 위해 본 연구에서는 모바일 기기를 이용하여 계량기 이미지를 획득하고, LTE 통신망을 이용하여 서버로 전송한 후, 광학 문자 인식 기술을 적용하여 획득한 이미지에서 사용량을 빠르게 인식하고 사용량 측정의 운영 효율을 향상시킬 수 있는 기술 구조를 제안하고자 한다.

II. 관련 연구

신경망 기술은 Rosenblatt (1958)에 의해 처음 제안되었으며, 신경망의 파라미터는 경험적 검색 방식을 통해 추정되었다. LeCun, Bottou, Bengio, and Haffner (1998)은 인간의 시각 지능 형성 과정을 모방한 합성곱 신경망 기술을 처음 제시하였고, Krizhevsky (2012)는 이를 발전시켜 합성곱 신경망의 층 수를 약 700여 개 적층하여 이미지 분류 분야에 적용함으로써 심층 신경망의 초기 구조를 제시하였다. 이후, 심층 신경망은 객체 검출, 객체 분할, 자제 추정 등 다양한 영상 분석 분야로 확대 적용되었다.

광학 문자 인식을 위해서는 문자열 위치 정보를 추정한 후 문자열을 인식하여야 한다.

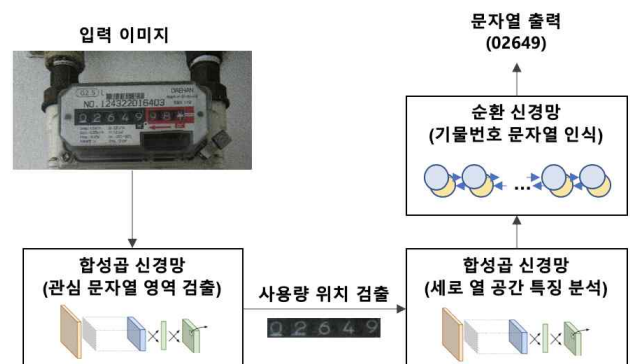
광학 문자 인식은 크게 2 가지 기술적 접근이 가능하다. 첫째는 단위 문자 각각의 위치 정보를 추정하고, 개별 문자 위치의 국지적 이미지를 분류 기술을 이용하여 인식하는 것이다 (LeCun et al., 1998). 이 방법은 개별 문자 각각을 검출하고 인식해야 하므로, 문자열을 구성하는 개별 문자 중 하나라도 인식 오류가 발생하면 문자열 전체의 인식 오류가 발생하는 문제점이 있다. 둘째는 복수개의 개별 문자 위치를 한 번에 검출하고 인식하는 방법이다 (Jaderberg, 2016; Liao, Shi, Bai, Wang, & Liu, 2017).

III. 연구 방법

계량기의 문자열 위치 정보를 추정하기 위해서 심층 신경망을 이용한 객체 검출 기술을 적용하였다. 객체 검출 기술은 두 가지 기술 구조로 나뉜다. 첫째는 객체 후보 영역을 먼저 찾고, 후보 영역의 객체 종류를 판별하는 2 단계 심층 신경망 구조이다 (Ren, 2015). 이 방법은 객체 후보 영역을 제안하는 영역 제안 심층 신경망을 별도로 두고 있어, 추가적인 연산량이 증가하는 단점이 있다. 둘째는 객체 영역 추정 및 객체의 종류를 한 번에 판별하는 1 단계 심층 신경망 구조이다 (Liu et al., 2016; Redmon, 2018). 이 방법은 이미지를 기 정의된 기저 사각형으로 분할하고, 각 기저 사각형의 좌표 정보를 변환하여 객체의 위치 정보를 추정하고, 객체 종류를 판별하는 구조로 이루어져 있으며, 2 단계 2 단계 기술 구조에 비해 연산량이 적은 장점이 있다. 본 연구에서는 1단계 심층 신경망 구조를 이용하여 문자열 위치 정보를 추정하였다.

문자열 영역 위치에 존재하는 문자열을 인식은 공간 정보를 순차 정보로 변환하고, 순차 정보를 순환 신경망을 이용하여 문자열로 변환하는 CRNN 기술을 적용하였다 (Shi, Bai, & Yao, 2017). CRNN은 문자열 영역을 합성곱 신경망을 거쳐 특징 맵으로 변환하고, 특징 맵의 세로 열 벡터를 순환 신경망으로 입력하여 문자열로 변환한다.

<그림 1> 문자 인식 기술 구조



IV. 실험 결과

<표 1>에 가스 사용량 인식 결과를 나타내었다.

<표 1> 가스 사용량 인식 결과.

데이터 분류		사용량 문자열 위치 추정			문자열 인식
		IoU	Recall	Prec.	
Normal		0.744	0.981	0.756	0.864
Level_1	Noise	0.722	0.969	0.733	0.805
	Reflex	0.672	0.979	0.679	0.757
	Scale	0.691	0.983	0.701	0.743
	Slant	0.519	0.993	0.522	0.615
Level_2	Noise	0.662	0.982	0.669	0.746
	Reflex	0.683	0.963	0.694	0.625
	Scale	0.652	0.982	0.664	0.717
	Slant	0.353	0.985	0.354	0.161

실험에 쓰인 데이터는 가스 계량기 이미지 27,120개 이다. 이 중 22,98개를 심층 신경망 학습 및 검증에 사용하였고, 4,135개를 테스트에 사용하였다. 테스트 이미지 중 멀리서 촬영하여 문자 영역 크기가 작거나, 이미지가 기울어졌거나, 잡음이 있거나, 빛반사가 있는 것을 scale, slant, noise 및 reflex로 구분하고, 통칭하여 abnormal 이미지로 지정하였다. Abnormal 이미지는 왜곡의 정도에 따라 정도가 심하지 않은 level_1 및 level_2로 분류하였다. Slant 데이터는 5 ~ 10도 기울어진 경우를 level_1, 10도 이상이면 level_2로 분류하였다. Scale 데이터는 문자열이 차지하는 영역이 전체 중 20 ~ 25%를 차지하면 level_1으로, 20% 이하이면 level_2로 분류하였다. 성능 측정은 위치 추정 및 문자열 인식에 대한 정확도를 측정하였고, 위치 추정에는 precision, recall 및 IoU (Intersection of Union)를 기준으로 측정하였다.

III. 참고문헌

Jaderberg, M., et al. (2016). Reading Text in the Wild with Convolutional Neural Networks. International Journal of Computer Vision, vol. 116, no. 1, 2016, pp. 1-20.

Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). Imagenet classification with deep convolutional neural networks. In Advances in neural information processing systems (pp. 1097-1105).

LeCun, Y., Bottou, L., Bengio, Y., & Haffner, P. (1998). Gradient-based learning applied to document recognition. Proceedings of the IEEE, 86(11), 2278-2324.

Liao, M., Shi, B., Bai, X., Wang, X., & Liu, W. (2017). Textboxes: A fast text detector with a single deep neural network. Paper presented at the Thirty-First AAAI Conference on Artificial Intelligence.

Liu, W., Anguelov, D., Erhan, D., Szegedy, C., Reed, S., Fu, C.-Y., & Berg, A. C. (2016). Ssd: Single shot multibox detector. Paper presented at the

European conference on computer vision.

Redmon, J., & Farhadi, A. (2018). Yolo3: An incremental improvement. arXiv preprint arXiv:1804.02767.

Ren, S., He, K., Girshick, R., & Sun, J. (2015). Faster r-cnn: Towards real-time object detection with region proposal networks. In Advances in neural information processing systems (pp. 91-99).

Rosenblatt, F. (1958). The perceptron: a probabilistic model for information storage and organization in the brain. Psychological review, 65(6), 386.

Shi, B., Bai, X., & Yao, C. (2017). An end-to-end trainable neural network for image-based sequence recognition and its application to scene text recognition. IEEE transactions on pattern analysis and machine intelligence, 39(11), 2298-2304.

E1.3 언어 모델링을 활용한 도메인 특화 신경망 기계 번역

이규복
연세대학교 산업공학과
glee48@yonsei.ac.kr

김우주
연세대학교 산업공학과
wkim@yonsei.ac.kr

Abstract - 정보통신기술의 발달로 서로 다른 언어를 구사하는 사람들 간에 일상 대화뿐만 아니라 특정 도메인의 지식도 공유할 기회가 유례없이 많아지고 있다. 일상 대화의 경우 구글 번역기와 같은 범용 번역기를 사용함으로써 간단한 의사소통이 가능해졌지만, 의료 및 법률과 같은 전문 분야의 지식은 범용 번역기의 사용만으로는 정확한 의미의 파악과 전달이 어렵다. 만약 특정 도메인에 특화된 번역기를 기업이나 개인이 직접 개발하고자 해도, 도메인에 특화된 병렬 데이터의 양이 많지 않아 새로운 기계 번역 모델 개발에 많은 제약 사항이 존재한다. 본 연구는 특정 도메인의 단일 언어 코퍼스의 경우 상대적으로 많은 양의 코퍼스를 확보할 수 있다는 점에 착안하여, 단일 언어 모델을 활용한 도메인 특화 신경망 기계 번역의 성능 향상 방법을 제안한다. 본 방법론은 저용량 데이터를 활용한 번역에서 흔히 사용되는 Back-translation 기법을 함께 사용할 수 있어 추가적인 성능 향상을 기대할 수 있는 장점을 가진다.

Key Terms - Domain-specific neural machine translation, Language model, Low-resource neural machine translation, Monolingual data

I. 서론

신경망 기계 번역(Neural Machine Translation; NMT)은 Sequence to sequence 모델의 등장으로 기계 번역 분야에서 눈부신 발전을 보여주었다(Bahdanau et al., 2014; Cho et al., 2014; Sutskever et al., 2014; Gehring et al., 2016). 특히 2017년 Google에서 개발한 Transformer(Vaswani et al., 2017) 구조는 기계 번역 분야뿐만 아니라 GPT(Radford et al., 2018)와 BERT(Devlin et al., 2018) 등의 사전 학습(pre-training) 모델의 형태로 여러 자연어 처리 분야에서 높은 성능을 달성할 수 있게 지대한 영향을 주었다.

높은 성능의 딥러닝 모델은 모델 구조의 비약적인 발전도 있지만, 대용량 데이터의 보유가 선행되어야 한다. 특히 기계 번역에서는 뉴스나 일상 구어체 데이터의 수집으로 도메인 비종속적(domain-agnostic)인 번역기의 개발은 어렵지 않게 이루어지고 있지만, 전문 분야와 같이 특정 도메인에 특화(domain-specific)된 번역기 개발은 여전히 심각한 데이터 부족 문제를 가지고 있다(Chu et al., 2018).

기존 도메인 특화 번역은 특정 도메인(in-domain)의 병렬 데이터와 비도메인(out-of-domain) 병렬 데이터를 함께 사용하여 학습하는 연구가 주류를 이룬다. Cost-weighting(Chen et al., 2017)은 도메인 이진 분류기(binary classifier)를 인코더와

연결하여 소스(source) 언어 문장을 기준으로 병렬 코퍼스의 도메인을 분류하게 우선 구축하고, 번역 모델의 손실 함수(loss function)에 도메인 확률을 반영하여 학습한다. 그 외 비도메인 번역기를 활용하는 방법으로는 naïve fine-tuning 방법과 mixed fine-tuning(Chu et al., 2017) 등이 있다.

도메인 특화 번역은 단일 언어(monolingual) 코퍼스를 활용하는 기계 번역 방법론을 적용하여 데이터 부족 문제를 완화할 수도 있다. 대표적인 방법론으로는 역방향 번역기를 활용하는 Back-translation(Sennrich et al., 2015)이 있다. 적은 양의 병렬 데이터로 역방향 번역기를 먼저 학습하고, 타겟(target) 언어의 단일 코퍼스를 역으로 번역하여 생성한 합성 소스(synthetic source) 문장을 구축한다. 이후, 기존 타겟 코퍼스와 합성 코퍼스를 페어링(pairing)하여 합성 병렬 데이터를 구축 후 기존 병렬 데이터와 동시에 데이터 셋으로 구축하여 정방향 번역기를 학습한다. 그 외 단일 언어 코퍼스를 활용하는 방법으로는 번역 작업과 언어 모델링 작업(task)을 동시에 진행하는 다중 작업 학습(Multitask Learning)법이 있다. 일반적으로 번역에서 사용하는 인코더-디코더 아키텍처(encoder-decoder architecture)의 쓰임은 소스 언어를 받아서 타겟 언어를 잘 예측하도록 구축하는 것인데, 인코더나 디코더는 개별적으로 한가지 언어의 임베딩(embedding)을 배우는 독립적인 언어 모델의 성격을 가질 수 있다. 이를 활용하여 인코더(Gulcehre et al., 2015; Zhang et al., 2016)와 디코더(Domhan et al., 2017)를 언어 모델(language model)로 개조하여 기존 번역 모델과 함께 학습하는 선행 연구들이 있다.

본 연구에서는 Transformer를 기반으로 기존 인코더-디코더 모델을 번역에 활용함과 동시에, 디코더가 언어 모델의 역할을 하는 다중 작업 학습법을 제안한다. 구체적으로는 Transformer 디코더의 인코더-디코더 어텐션(encoder-decoder attention) 전후에 스킵 커넥션(skip connection)을 삽입함으로써, 번역 작업은 기존 모델의 데이터 흐름대로, 언어 모델링은 스킵 커넥션을 사용하여 순차적으로 학습한다. 특히 언어 모델링 학습 시 인코더와 인코더-디코더 어텐션의 파라미터를 동결(freeze)시켜 각 디코더 레이어의 역할이 구분되어 학습되기에 스킵 커넥션을 제거하면 일반적인 Transformer 모델과 동일하게 된다. 이러한 점 때문에 본 방법론은 Back-translation 기법을 함께 사용할 수 있어 추가적 성능 향상을 기대할 수 있다는 장점을 가진다.

II. 참고문헌

- Bahdanau, D., Cho, K., & Bengio, Y. (2014). Neural machine translation by jointly learning to align and translate. arXiv preprint arXiv:1409.0473.
- Chen, B., Cherry, C., Foster, G., & Larkin, S. (2017, August). Cost weighting for neural machine translation domain adaptation. In *Proceedings of the First Workshop on Neural Machine Translation* (pp. 40-46).
- Cho, K., Van Merriënboer, B., Gulcehre, C., Bahdanau, D., Bougares, F., Schwenk, H., & Bengio, Y. (2014). Learning phrase representations using RNN encoder-decoder for statistical machine translation. arXiv preprint arXiv:1406.1078.
- Chu, C., Dabre, R., & Kurohashi, S. (2017). An empirical comparison of simple domain adaptation methods for neural machine translation. arXiv preprint arXiv:1701.03214.
- Chu, C., & Wang, R. (2018). A survey of domain adaptation for neural machine translation. arXiv preprint arXiv:1806.00258.
- Devlin, J., Chang, M. W., Lee, K., & Toutanova, K. (2018). Bert: Pre-training of deep bidirectional transformers for language understanding. arXiv preprint arXiv:1810.04805.
- Domhan, T., & Hieber, F. (2017, September). Using target-side monolingual data for neural machine translation through multi-task learning. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing* (pp. 1500-1505).
- Gehring, J., Auli, M., Grangier, D., & Dauphin, Y. N. (2016). A convolutional encoder model for neural machine translation. arXiv preprint arXiv:1611.02344.
- Gulcehre, C., Firat, O., Xu, K., Cho, K., Barrault, L., Lin, H. C., ... & Bengio, Y. (2015). On using monolingual corpora in neural machine translation. arXiv preprint arXiv:1503.03535.
- Radford, A., Narasimhan, K., Salimans, T., & Sutskever, I. (2018). Improving language understanding by generative pre-training. URL <https://s3-us-west-2.amazonaws.com/openaiassets/researchcovers/languageunsupervised/languageunderstandingpaper.pdf>.
- Sennrich, R., Haddow, B., & Birch, A. (2015). Improving neural machine translation models with monolingual data. arXiv preprint arXiv:1511.06709.
- Sutskever, I., Vinyals, O., & Le, Q. V. (2014). Sequence to sequence learning with neural networks. In *Advances in neural information processing systems* (pp. 3104-3112).
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. In *Advances in neural information processing systems* (pp. 5998-6008).
- Zhang, J., & Zong, C. (2016). Exploiting source-side monolingual data in neural machine translation.

E1.4 정보보호 대책의 성능을 고려한 투자 포트폴리오의 게임 이론적 최적화

Game Theoretic Optimization of Investment Portfolio Considering the Performance of Information Security Countermeasure

이상훈 (Sang-Hoon Lee)

소속(국문): 충북대학교 경영정보학과

소속(영문): Department of MIS, Chungbuk National University

E-mail: sanghoon@chungbuk.ac.kr

김태성 (Tae-Sung Kim)

소속(국문): 충북대학교 경영정보학과

소속(영문): Department of MIS, Chungbuk National University

E-mail: kimts@chungbuk.ac.kr

Corresponding author: Tae-Sung Kim

1 Chungdae-ro, Seowon-gu, Cheongju-si, Chungbuk 28644 South Korea

E-mail: kimts@chungbuk.ac.kr

국문초록

사물 인터넷, 빅데이터, 클라우드, 인공지능 등 다양한 정보통신기술이 발전하고 있고, 정보보호의 대상이 증가하고있다. 정보통신기술의 발전에 비례해서 정보보호의 필요성이 확대되고 있지만, 정보보호 투자에 대한 관심은 저조한 상황이다. 또한 정보보호 대책의 종류와 특성이 다양하기 때문에 객관적인 비교와 평가가 힘들고, 객관적인 의사결정 방법이 부족한 실정이다. 이에 본 연구에서는 게임 이론을 이용하여 정보보호 대책 투자 포트폴리오를 구성하는 방법을 제안하고 선형계획법을 이용하여 최적 방어 확률을 도출한다. 2인 게임 모델을 이용하여 정보보호 담당자와 공격자를 게임의 경기자로 구성한 뒤, 정보보호 대책을 정보보호 담당자의 전략으로, 정보보호 위협을 공격자의 전략으로 각각 설정한다. 다수의 정보보호 위협에 따른 정보보호 대책이 배치된 환경에서 정보보호 대책의 방어 비율을 이용하여 정보보호 대책 투자 포트폴리오를 산출한다. 최적화된 포트폴리오를 이용하여 방어 확률을 최대화하는 게임 값을 도출한다. 정보보호 대책의 실제 성능 데이터를 이용하여 수치 예제를 구성하고, 제안한 게임 모델을 적용하고 평가한다. 조직의 정보보호 담당자는 본 연구에서 제안한 방법을 이용하여 정보보호 대책의 성능을 고려한 효과적인 정보보호 대책 투자 포트폴리오를 수립할 수 있을 것이다.

주제어

게임이론, 정보보호, 투자 포트폴리오

E1.5 랜덤성에 기반한 인공지능 오목의 구현

남건민
서울과학기술대학교 빅데이터경영공학
gibicee@naver.com

천세학
서울과학기술대학교 빅데이터경영공학
shchun@seoultech.ac.kr

Abstract - 알파고(Alpha)의 등장은 전 세계에서 AI(인공지능)에 대한 관심을 폭발적으로 증가시켰다. 인간만의 전유물이라고 여겨지던 바둑에서 알파고는 딥러닝과 강화학습을 통해 팔목할 만한 성능을 보여주었다. 이를 통해 AI가 인간을 얼마나 모방할 수 있는가에 대한 가능성이 크게 대두되었다. 본 논문은 알파고에 사용된 방법론들을 적용하여 파이썬(Python)을 통해 오목 AI를 구현해보자 한다.

Key Terms - 강화학습, 딥러닝, 알파고, 오목, 인공지능

I. 서론

바둑에서 최선의 수를 찾고자 하는 시도는 예전부터 있었다. 하지만 바둑의 경우의 수는 약 10^{170} 으로, 우주의 원자 수보다 많다는 사실 때문에 바둑의 최선의 수를 컴퓨터를 통해 계산한다는 것은 불가능하다고 여겨졌다. 그래서 바둑은 인간만이 가능한 영역이라고 불리기도 했다. 하지만 딥러닝을 활용한 알파고의 등장 이후, 그러한 개념은 파괴되었다. 이제는 대부분의 바둑 프로 기사들이 AI로 훈련하고 상당한 도움을 받는다고 한다. 이러한 인공지능 분야의 비약적인 발전은, AI가 인간의 직관이나 직감을 모방할 수 있다는 가능성을 보여준다. 앞으로 기술이 더욱 발전한다면 인공지능을 통해 좀 더 복잡한 문제를 해결할 수 있을 것이란 잠재력이 있다.

본 논문은 바둑보다는 경우의 수가 적은 오목 환경에서 동작하는 AI를 구현하는 것이지만, 인공지능의 가능성을 직접 확인할 수 있다는 것에 의미를 둘 수 있다.

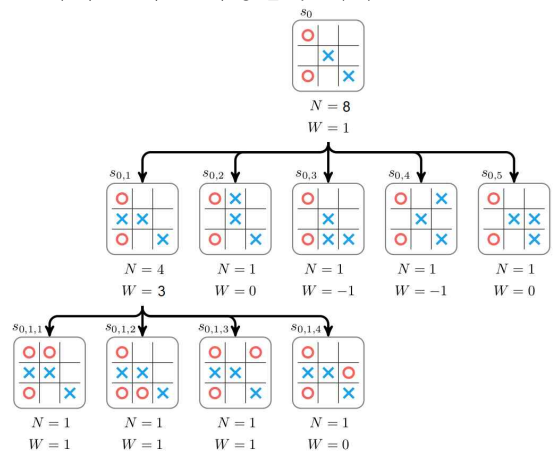
II. 기존에 사용된 방법론

알파고에 사용된 방법론을 요약하자면 다음과 같다. 알파고는 이미 존재하는 바둑 기사의 게임 데이터로부터 약 3천만 수를 가져와 정책망(Policy Network)을 학습한다. 그렇게 지도 학습(Supervised Learning)된 다양한 버전의 정책망끼리 서로 대전하고 스스로의 대전 기록을 복기하며 강화학습을 통해 정책망의 성능을 개선한다. 이렇게 학습된 정책망은 이후 MCTS(몬테카를로 트리 탐색) 알고리즘에서 활용된다^[1].

알파고 제로(AlphaGo Zero)는 알파고 이후에 공개된 버전이다. 알파고 제로와 알파고의 가장 큰 차이점은 학습 과정에 있다. 알파고는 초기에 인간이 플레이한 기록을 토대로 학습을 시작한 반면, 알파제로는 그러한 인간의 플레이 데이터를 전혀 사용하지 않고 학습한다. 오직 바둑의 규칙만을 알려주어 그 규칙을 토대로 스스로 대국하여 학습을 진행한다. 학습 초기에는 무작위로 돌을 두었지만 학습이 어느 정도

진행된 이후 알파고 제로는 알파고보다 훨씬 뛰어난 성능을 보여주었다고 한다^[2].

MCTS(몬테카를로 트리 탐색) 알고리즘은 다양한 경우의 수를 탐색하여 최종적으로 착수 위치를 결정하는 알고리즘이다. 3x3 보드에서 같은 글자로 한 줄을 만들어야 하는 틱택토(Tic Tac Toe) 게임을 예시로 하여 설명하겠다. 아래 그림은 틱택토 게임에서 특정 상태를 루트(root) 노드로 하여 트리를 작성한 것이다.



<그림 1> 몬테카를로 트리 탐색 예시

그림1의 트리는 가장 위에 있는 루트 노드인 상태 S_0 에서 플레이어 X가 착수할 수 있는 5가지 경우에 따른 자식 노드로의 확장을 보여준다. 노드의 방문횟수를 N, 기댓값을 W로 했을 때, 상태 $S_{0,1}$ 으로 4번 확장하여 3번 승리했음을 알 수 있다. 만약 특정 상태를 방문한 경우에서 패배하면 W값에 -1이 적용되고, 반대로 승리하면 W값에 +1이 적용된다. 이렇게 특정 상태의 N값과 W값에 대한 정보가 갱신되면, 이를 자식 노드로 가지는 부모 노드의 정보 또한 갱신되는 역전파(backpropagation) 과정을 수행하게 된다. 위 트리를 참고했을 때, 현재 게임의 상태가 S_0 라면 플레이어 X는 가장 큰 기댓값(W)을 가지는 상태 $S_{0,1}$ 가 되도록 돌을 착수할 것이라고 예측할 수 있다.

하지만 모든 경우의 수를 트리 탐색하는 것은 많은 시간을 필요로 하기 때문에 중요한 위치(착수할 확률이 높은 위치)는 다른 위치보다 높은 비중으로 탐색하는 알고리즘이 필요하다. 이를 MCTS 알고리즘이라고 한다. 이 과정에서 중요한 위치가 어디인지 파악하는 것에 학습된 정책망이 활용된다.

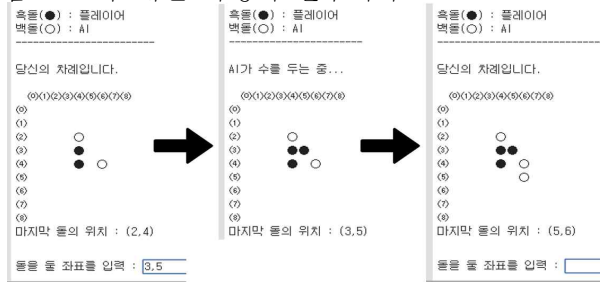
III. 파이썬을 이용한 인공지능 오목의 구현²⁾

2)) 구현에 사용된 소스코드 및 AI와 대결하는 방법 등은

다양한 딥러닝 라이브러리를 지원하는 파이썬(Python)과 행렬과 같은 수치 계산에 유용한 Numpy, Theano 등의 라이브러리를 사용하여 오목 AI를 구현하였다. 오목 게임의 기본적인 규칙은 렌주(Renju)를 기반으로 한다. 렌주는 오목에서 흑백을 공평하게 하기 위해 만들어진 규칙이다. 흑만 3-3과 4-4 그리고 장목이 금지된다. 렌주에서는 15x15 크기의 보드를 사용하지만, 학습에 필요한 시간과 인공지능의 기대 성능에 문제가 있어서 9x9 보드로 크기를 축소하였다.

학습 과정에서 인간의 게임 플레이 데이터를 사용하지 않는 알파고 제로의 방법론을 적용하였다. 초기의 정책망은 완전한 무작위로 돌을 두지만, 반복된 학습을 통해 정책망의 성능이 개선된다. 정책망의 성능이 개선될수록 AI는 오목판에서 중요한 위치가 어디인지 잘 파악할 수 있게 된다. 그로 인해 좀 더 효율적인 MCTS 알고리즘을 수행할 수 있고 이것이 AI의 성능에 직접적인 영향을 준다.

플레이 과정에서는 매 턴마다 오목판의 상태를 출력하며, 플레이어(사람)의 턴에서는 착수 위치를 입력 받는다. 이후 번갈아 가며 AI가 착수한다. 아래 그림 2는 Google Colab 환경에서의 오목 대결 과정의 일부를 보여주고 있다. 플레이어(사람)와 AI 오목 간의 게임 화면으로 셀을 실행하면 오목판의 상태를 출력하며, 플레이어로부터 착수 위치를 입력 받는다. 이후 번갈아 가며 AI가 착수한다. 아래 그림은 오목 대결 과정의 일부이다.



<그림 2> 플레이어와 오목 AI의 대결

학습 횟수에 따른 오목 AI의 성능에 대한 주관적인 의견은 다음과 같다.

<표 1> 학습 횟수에 따른 인공지능의 성능.

학습횟수	성능에 대한 주관적 의견
1000	거의 무작위로 돌을 두는 것처럼 보임.
2000	자신의 돌을 3~4목으로 이으려고 함.
3000	수비보다 공격을 중요시하기 때문에
4000	수비해야 할 상황에서 공격을 하는 경우가
5000	적음.
6000	수비를 중요시하기 시작함. 그러나 간혹
7000	이상한 실수를 하기도 함.
8000	집중해서 플레이하지 않으면 이기기
9000	힘들어질 정도로 성능이 개선됨.
10000	

IV. 결론

학습을 반복할수록 오목 AI의 성능이 개선되는 것을 관찰할 수 있었다. 이를 통해 인간의 플레이 데이터가 없어도 랜덤성에 기반하여 강화학습을

적용한 결과, 인공지능의 구현이 가능했다는 것을 확인할 수 있었다.

V. 참고문헌

- [1] Silver, D., Huang, A., Maddison, C. et al. Mastering the game of Go with deep neural networks and tree search. Nature 529, 484–489 (2016).
- [2] Silver, D., Schrittwieser, J., Simonyan, K. et al. Mastering the game of Go without human knowledge. Nature 550, 354–359 (2017).
- [3] 박현수, 김경중. (2013). 게임 인공지능 최신 연구 동향. 정보과학회지, 31(7), 8-15.
- [4] 위키백과, “알파고”, <https://ko.wikipedia.org/wiki/알파고>
- [5] 브런치, "AlphaGo에 적용된 딥러닝 알고리즘 분석", <https://brunch.co.kr/@justinleeanac/2>
- [6] 티스토리 블로그, "알파고 ZERO: AlphaGo Zero - How and Why it works 포스팅 번역", <https://leekh7411.tistory.com/1?category=768501>

A2 [기획세션] 디지털 비즈니스 I

A2.1 Deep Learning기반 지능형 CCTV

조명돌 - (주)마크애니 최고경영관리본부장

A2.2 포스트 팬데믹 시대의 AI학습의 중요성

박서린 - (주)N3N CLOUD 매니저

B2 [특별세션] 인텔리전스 대상 기업세션 I

B2.1 uTH 2.0 - 지능형 데이터 기반 디지털무역물류 플랫폼

김채미 - (주)한국무역정보통신 디지털무역본부장

B2.2 인공지능 챗봇 소개 및 활용 사례

장정훈 - (주)와이즈넷 성장기술연구소장

B2.3 Hyper Inference - 실사수준 렌더링 활용 객체탐지 학습 솔루션 (기아자동차 사례중심)

조만희 - (주)메가존 구글클라우드 테크팀

B2.4 빅데이터 기반 글로벌 케이팝 플랫폼 "후즈팬"

곽영호 - (주)한터글로벌 대표이사

C2 [특별세션] 인공지능 응용 경진대회 I

C2.1 시계열 클러스터링을 이용한 이터닝 학습자 유형 변화 분석

임종언, 김민지, 남윤주, 류예나, 이채연(국민대)

C2.2 중첩적 클러스터링을 활용한 군집별 서비스 제안

김민수, 김효중, 신동훈(연세대)

C2.3 학습행동에 따른 학생 군집별 중도해지예측모델

이재민, 이호진, 김재우, Janine Anne Laddaran, Liu Zih Syuan(연세대)

D2 [학술논문] 코로나 19 팬데믹

D2.1 오토인코더 기반 특징추출을 통한 국가별 COVID-19 일일 신규 확진자 수 예측 모델링

홍성원
충북대학교 경영학부
blockdeal@naver.com

Abstract - 본 연구에서는 COVID-19의 국가별 신규 확진자 수를 나타내는 원천 데이터만을 활용한 AE와 RNN 기반의 예측 모델링 접근법을 제시한다. 이는 AE로부터 추출된 특징으로 RNN을 훈련시켜 신규 확진자 수를 예측하는 접근법으로, 길이가 짧은 원천 데이터를 활용한 모델링에 유용하다. 중국, 미국, 브라질 3개국의 일일 신규 확진자 수 예측에 적용한 결과 예측력이 개선되었다.

Key Terms - AE(autoencoder), LSTM(long short-term memory), RNN(recurrent neural network), 전염병 예측 모델링, 특징추출

I. 서론

COVID-19 팬데믹을 조기 극복하기 위한 시도의 하나로 각국 정부는 외국인 출입국을 통제하고 있다. 하지만 이에 따라 팬데믹 경제위기가 가중될 위험이 있기 때문에 일부 국가에서는 외국인 출입국 통제를 완화해야 한다는 목소리가 커지고 있으며, 구체적으로 어느 시점에 완화할 것인지에 관해 활발한 논의가 이루어지고 있다. 또한 팬데믹으로부터 직접적인 타격을 입은 수출기업들도 향후의 경영환경 변화에 민첩하게 대응하기 위해 주요 국가들의 팬데믹 현황을 예의주시하고 있으며, 경영정상화를 위해 각종 데이터 분석 역량을 동원하고 있다. 이처럼 국가별 신규 확진자 수 예측은 출입국 통제 완화의 적정 시기를 결정 및 추정하는 데 활용 가치가 있다.

국가별 신규 확진자 수의 변동에는 COVID-19 자체의 확산세뿐만 아니라, 각국의 상이한 의료 체계와 유행인구율, 사회적 거리두기에 대한 인식 등의 수많은 변인들이 영향을 미치는 것으로 볼 수 있다. 하지만 정보통신기술 인프라가 발달하지 않은 개발도상국가들에 대한 데이터 확보가 용이하지 않으며 정량화 자체가 어려운 변인도 적지 않기 때문에 전염병 예측 모델링에 어려움이 있다. 이를 개선하기 위해 본 연구에서는 국가별 일일 확진자 수의 추이를 나타내는 원천 데이터와 AE 및 RNN 두 가지 신경망을 활용한 전염병 예측 모델링 접근법을 제시한다.

II. 특징추출 기반 순차적 모델링 접근법

일반적으로 AE는 입력층부터 중간의 코딩층으로 갈수록 노드의 개수가 줄어드는 구조를 가지면서 입력과 출력이 같아지도록 훈련된 다층 신경망으로부터 얻어지며, 차원축소에 있어 전통적 기법보다 양호한 성능을 기대할 수 있는 비지도 학습 모델이다(Hinton & Salakhutdinov, 2006). 차원의 저주 문제를 개선하는 보통의 AE에 의해 축소된 특징 데이터는 축소 이전의 입력 데이터보다 저차원을 가지면서 입력 데이터를 잘 설명하는 특징들을 내포하고 있다. AE는 이처럼 차원축소 및 특징추출을 위해 널리 활용되며, 차원축소의 효과를 높이기 위해 코딩 노드를 규제하는 SAE(sparse

autoencoder) 등의 여러 변형 모델이 제시되어 왔다(Makhzani & Frey, 2013).

이러한 AE의 차원축소 기능은 일반적인 순차적 모델링 문제에서도 효과적이지만, COVID-19 일일 신규 확진자 수의 경우와 같이 시계열의 길이가 짧고 목표값에 영향을 미치는 요인에 대한 프로파일링이 어려운 경우에는 데이터 크기 자체가 작기 때문에 차원축소 접근법이 부적합하다. 본 연구에서 제시하는 접근법은 보통의 AE와는 반대로 코딩층의 노드 개수를 입력층보다 크게 해 차원축소의 반대 효과를 얻는 방법으로 RNN 기반의 순차적 모델링 시 데이터 크기 증가에 따른 성능 개선을 기대할 수 있다.

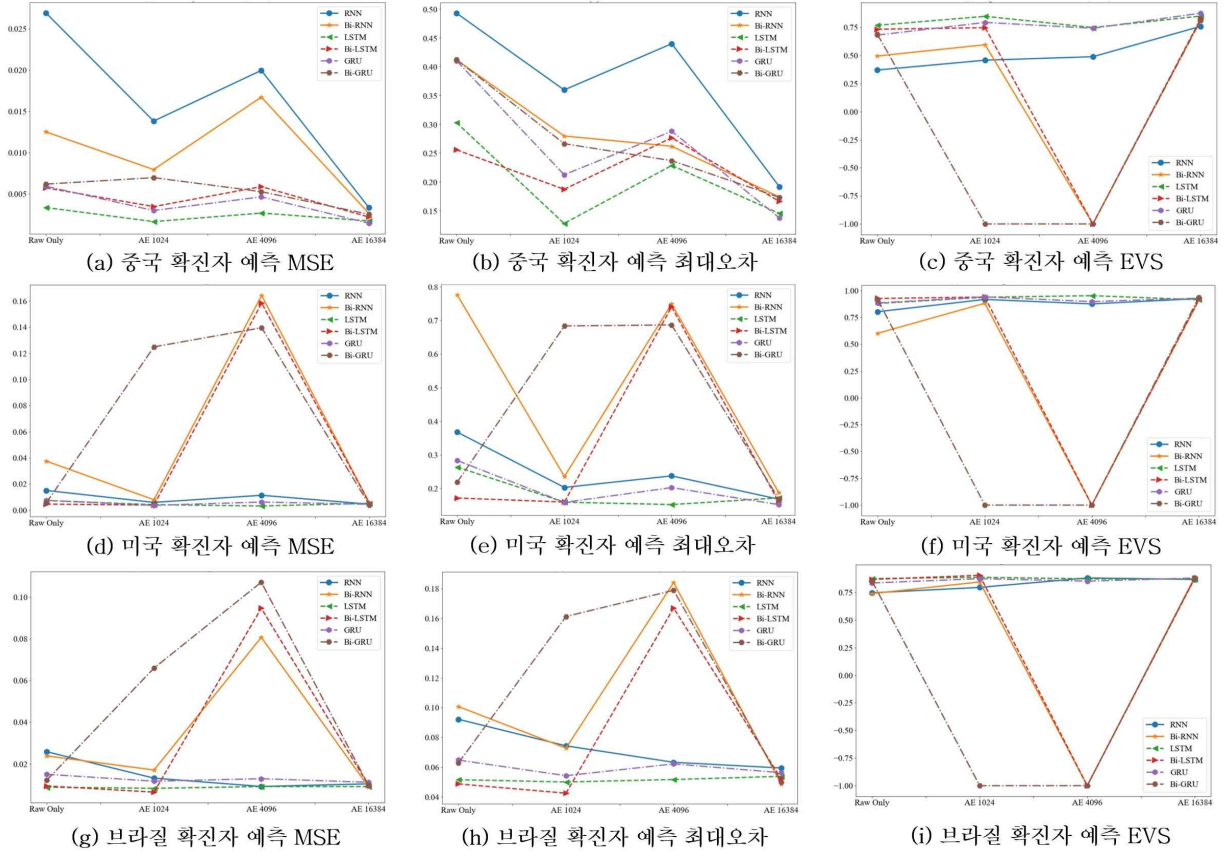
동 접근법의 예측력 개선 여부를 확인하기 위해 2020년 1월 1일부터 6월 11일까지 전세계 195개국의 일일 신규 확진자 수 데이터에 적용한다. 예측 대상은 중국, 미국, 브라질의 COVID-19 신규 확진자 수이며, [0,1] 범위로 전처리된 시계열 데이터셋을 9:1로 분할해 훈련 및 테스트셋을 구성한다.

AE는 코딩층이 각각 1024, 4096, 16384개의 노드로 구성된 세 개의 모델을 활용하며, 순차적 모델링을 위한 RNN 모델로는 단순(vanilla) RNN, Bi-RNN, LSTM, Bi-LSTM, GRU, Bi-GRU 여섯 종류의 성능을 비교 검토한다(Goodfellow et al., 2016). 이에 따라 훈련 및 테스트 시계열의 길이는 각각 130, 33이며, 국가별 확진자 수 예측에서 성능의 비교 대상이 되는 모델의 개수는 각각 24개이다.

III. 예측 결과 및 논의

제안 접근법을 통한 예측 결과, 먼저 중국의 경우 <그림 1>의 (a), (b), (c)에 나타난 바와 같이 MSE (mean squared error), 최대 오차(max error)의 측면에서 단순 LSTM이 평균적으로 Bi-LSTM을 비롯해 나머지 네 가지 모델들을 상회하는 성능을 보였고 EVS(explained variance score)가 0.75 수준으로 일정 수준 설명력이 있는 것으로 나타났다. RNN 및 Bi-RNN의 경우 원천 데이터만을 활용했을 때(Raw Only) 비교적 높은 오차를 보였으나 AE의 코딩 차원이 16,384가 됨에 따라 성능이 제고됐다. 미국의 경우에도 LSTM이 전반적으로 양호한 성능을 보였으며 EVS도 적절하게 나타났다. Bi-GRU 등의 여타 모델들은 AE의 코딩 차원이 4,096일 때 도리어 가장 악화된 성능을 보였으나, 코딩 차원이 16,384가 될 때 성능이 확연히 제고되었으며, 브라질의 경우도 이와 유사하게 나타났다. Bi-RNN, Bi-LSTM, Bi-GRU 등의 양방향 모델은 훈련가능한 파라미터의 수가 RNN, LSTM, GRU보다 많음에도 불구하고 중국, 미국, 브라질의 일일 신규 확진자 수 예측 모두에서 상대적으로 불안정한 성능을 보였다.

제안 접근법은 AE의 코딩 노드 개수를 입력 노



<그림 1> 중국, 미국, 브라질 일일 신규 확진자 수 예측 결과

드 개수보다 크게 해 특징을 추출한 후 이것으로 RNN 모델을 훈련시키는 기법으로 시계열의 길이가 짧고 데이터 프로파일링이 제한된 상황에서 활용될 수 있다. 제안 접근법의 개선 방향은 크게 두 가지이다. 첫째, AE의 은닉 노드를 제한 볼츠만 머신으로 구성해 보다 다양한 측면에서 여러 특징을 추출하는 것이다(Hinton & Salakhutdinov, 2006; Goodfellow et al., 2016). 둘째, COVID-19 일일 확진자 수 데이터로부터 신경망이 계층적 표상을 학습할 수 있다는 것을 확인한 본 연구의 개선, 보완으로 DAE(denoising autoencoder) 및 VAE(variational autoencoder)를 활용해 예측력이 개선될 수 있는지 검토하는 것이다.

IV. 참고문헌

Goodfellow, I., Y. Bengio, and A. Courville, *Deep Learning*, MIT press, 2016.

Hinton, G.E. and R.R. Salakhutdinov, "Reducing the dimensionality of data with neural networks," *science*, 313(5786), 2006.

Markhzani, A., and B. Frey, "K-sparse autoencoders," *arXiv preprint arXiv:1312.5663*, 2013.

D2.2 SIR 모형과 LSTM 기반의 COVID-19 확진자 예측 비교 및 분석

Comparison and Analysis of COVID-19 Confirmed Cases Based on the SIR model and LSTM

김진오 (Jino Kim)

소속(국문): 경북대학교 컴퓨터학부
소속(영문): School of CSE, Kyungpook University
E-mail: kimjino1996@gmail.com

김지수 (Jisu Kim)

소속(국문): 경북대학교 컴퓨터학부
소속(영문): School of CSE, Kyungpook University
E-mail: wlt4511@gmail.com

김희원 (Heewon Kim)

소속(국문): 경북대학교 컴퓨터학부
소속(영문): School of CSE, Kyungpook University
E-mail: kvlksgjp1@naver.com

박효상 (Hyosang Park)

소속(국문): 경북대학교 컴퓨터학부
소속(영문): School of CSE, Kyungpook University
E-mail: latter2005@gmail.com

김지혜 (Jihye Kim)

소속(국문): 경북대학교 컴퓨터학부
소속(영문): School of CSE, Kyungpook University
E-mail: jihye12.5@gmail.com

옥혜원 (Hyewon Ok)

소속(국문): 경북대학교 자연과학대
소속(영문): School of NS, Kyungpook University
E-mail: kkkko30@naver.com

문혜원 (Hyewon Moon)

소속(국문): 경북대학교 컴퓨터학부
소속(영문): School of CSE, Kyungpook University

정명훈 (Hyewon Moon)

소속(국문): 구글코리아
소속(영문): Google Korea
E-mail: javalove93@naver.com

고석주 (Hyewon Moon)

소속(국문): 경북대학교 컴퓨터학부
소속(영문): School of CSE, Kyungpook University
E-mail: sjkoh@knu.ac.kr

Corresponding author: Seokju Ko

Kyungpook Univ, Sangyeok 3-dong, Buk-gu, Daegu, Korea
Tel: +82-53-950-5550, Fax: +82-53-957-4846, E-mail: cse-admission@knu.ac.kr

국문초록

2019년 12월 중국 후베이성 우한 시에서 시작된 코로나바이러스감염증-19(이하 COVID-19)은 2020년 1월부터 전 세계로 퍼져, 일부 국가 및 지역을 제외한 대부분의 나라와 모든 대륙으로 확산되었다. 이에 따라 WHO는 2020년 1월 30일 국제적 공중보건 비상사태(PHEIC)를 선포한 후, 2020년 3월 11일 범 유행 전염병(Pandemic)을 선언하였다. COVID-19의 영향은 생명 손실에 국한되지 않고, 사회경제적 영역으로 확대되고 있다. 경제 침체로 인한 문제, 신체적, 물질적 피해와 지속적인 스트레스, 해소되지 않는 사회적 불안 등을 감소시키기 위해 전염병 확산에 대한 예측은 필수 불가결 하다.

본 논문은 이처럼 최근 유행하고 있는 전염병인 COVID-19에 대해 전염병 모형과 순환 신경망(Recurrent Neural Network) 두 기법을 통해 확산을 예측하고 결과를 비교한다. 전 세계의 국가별 확진자와 사망자 통계를 활용하여 예측 결과를 도출하였다. 전염병 모형을 통한 예측에는

SIR(Susceptible-Infected-Recovered) 모형을 사용하며, 기계 학습의 최적화 알고리즘 중 L- BFGS를 사용하여 SIR 모형의 곡선 적합을 진행하였다. RNN 기반 예측 모델로는 LSTM 모델을 사용하였다. 예측을 실시하여 각 기법의 결론을 도출하고 실제 확진자 통계와의 정확도를 비교하였다. 본 논문은 특정한 나라들에 대하여 전염병 초기의 SIR 모형을 사용한 예측과, 현재(2020.5.29) 까지 진행된 상태에서의 예측을 비교한다. 두 예측의 오차를 통하여 방역 수준에 대한 추론을 제시한다. 추가 확진자와 누적 확진자의 관계 그래프를 제공, 기술키와 두 예측 그래프 사이의 오차와 관계성을 확인하고 대략적 방역 수준을 정의한다.

주제어

SIR, COVID-19, RNN, 기계학습.

Acknowledgement (사사)

본 연구는 과학기술정보통신부 및 정보통신기획평가원의 SW중심대학사업의 연구결과로 수행되었음(2015-0-00912)

1. 서론

코로나바이러스감염증-19(이하 COVID-19)은 2019년 12월 중국 후베이성 우한 시에서 시작, 2020년 1월부터 전 세계로 퍼지기 시작하여, 일부 국가 및 지역을 제외한 대부분의 나라와 모든 대륙으로 확산되었다. 2020년 1월 30일 WHO에서 국제적 공중보건 비상사태(PHEIC)를 선포하고 잇따라, 3월 11일 범유행전염병(Pandemic)을 선언하였다.

현재 2020년 5월 기준 COVID-19는 세계적으로 400만 명 이상의 확진자와 30만 명 이상의 사망자를 기록하고 있다. 국내의 경우 2020년 1월 20일에 첫 확진자가 발생한 이후 대구-경북 지역을 중심으로 폭발적인 확산이 진행되었고, 2020년 5월 15일을 기준으로 11,018명의 확진자와 260명의 사망자가 보고되었다.

COVID-19의 영향은 생명 손실에 국한되지 않고, 사회경제적 영역으로 확대되었다. 전 세계적인 국경 폐쇄 및 여행 금지 조치에 의해 국가 간 국경의 벽이 높아졌으며, 경제 침체로 인한 문제, 신체·물질적 피해와 지속적인 스트레스, 해소되지 않는 사회적 불안 등이 발생하고 있다.

이러한 상황에서, 각 국가별 확산 억제를 위한 정책 및 계획의 수립에 우리가 만든 모델을 통한 확진자와 피크의 예측이 큰 도움을 줄 수 있다.

전염병에 대한 억제·대응 책의 모색은 기본적으로 질병 확산 모형(disease spread model)을 바탕으로 여러 대응책을 실험한 후 그 효과를 분석하는 것으로 이루어지며, 이러한 모형의 구축에는 요구되는 정확성, 시스템 규모 등에 따라 다양한 수학, 최적화, 시뮬레이션 기법들이 활용되고 있다.

본 논문에서는 확진자의 시계열 데이터와 SIR 모델을 활용하여 이후의 확진자 수와 최대 확진 지점을 예측하는 모델을 소개한다. 예측을 위한 데이터는 존스 홉킨스 대학의 시스템 사이언스 및

엔지니어링 센터(CSSE)에서 제공하는 집계 데이터를 사용하였으며, SIR 모형을 통한 예측의 경우 SIR 모델의 인자 값을 기계 학습을 통해 추정하는 방식을 이용해 구현하였다.

2. 본론

2.1 개요

COVID - 19의 확진자를 예측하기 위해서 LSTM 과 SIR 모형 기반 예측모형을 구현한다. 두 모델을 사용하여 예측한 결과 값을 비교하고, 장단점을 파악하여 사용자의 목적에 따라 적합한 모델을 선택할 수 있도록 한다. 또한, 미국, 러시아, 이스라엘 등의 나라를 대상으로 진행 정도 (~초기, ~중기) 별로 예측을 실시한다. 예측은 SIR 모형 기반 예측 모델로 진행한다. 해당 예측들의 오차를 통해, 각 나라 별 방역 수준을 추론하고, 추가확진자와 누적확진자의 관계그래프와 같이 확산 속도를 시각화한 자료와 비교한다.

1 예측모델 개발

2 LSTM 기반 예측 모델 개발

누적 확진자 수와 같은 시계열 데이터를 예측하는데 적합한 RNN 계열의 모델을 사용하였으며, RNN 모델의 장기 의존성 문제를 해결하기 위해 LSTM을 사용하였다.

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f)$$

각각의 날짜에 해당하는 확진자, 사망자, 회복자 숫자에 대한 데이터와 이들의 차분 값, 관계적 변화 값을 인풋 데이터로 하여 미래의 확진자 수를 예측한다. 예측된 결과는 추가적인 미래 확진자 수 예측의 인풋 데이터로 재사용 될 수 있다.

학습을 위한 활성화함수로는 ReLU를 사용하였으며

경사도 계산에 대해서는 Adam Optimizer를 사용하였다.

3 SIR 모형 기반 예측 모델 개발

SIR 모형은 특정 폐쇄 집단의 인구가 전염병에 감염되는 단계를 S(susceptible), I(infected), R(recovered)의 세 단계로 나누어, 시간에 따라 각 단계(S→I, I→R)로 이동하는 정도를 나타낸 전염병 모형이다.

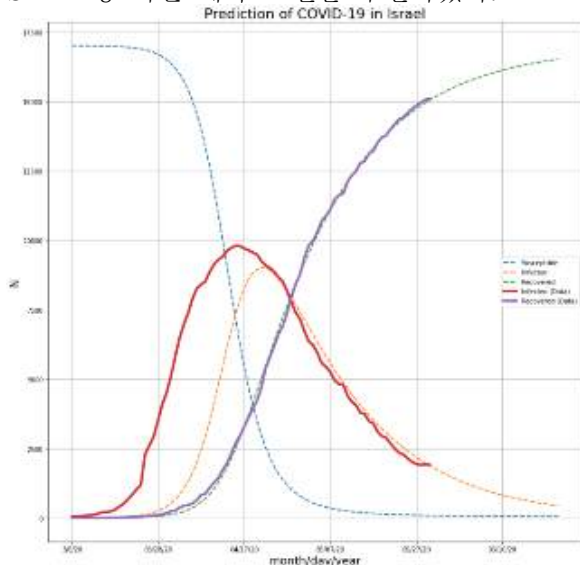
SIR 모형은 3 $\frac{dS}{dt} = -\beta SI$ 개의 비선형 미분 방정식으로 구성된다.

$$\frac{dI}{dt} = \beta SI - \gamma I$$

$$\frac{dR}{dt} = \gamma I$$

여기서 t 는 시간, $S(t)$ 는 감염 가능성이 있는 대상의 수, $I(t)$ 는 감염이 된 대상의 수, $R(t)$ 는 사망하거나 회복한 개체의 집단을 일컫는다. 이때 β 는 감염률(infection rate), γ 는 회복률(the recovery rate)을 나타낸다.

해당 모델을 기반으로, 전체 모집단의 개체 수인 N과 사용할 데이터 구간을 조절 가능하게 하여, 데이터에 적합한 SIR 모델에서의 β , γ 값을 찾는 SIR 모형 기반 예측 모델을 구현하였다.



<Figure 1>Example of SIR Prediction Model

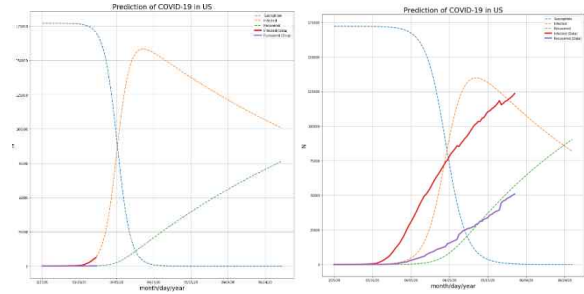
구현한 SIR 모형은 S, I, R의 단계에 있는 개체의 수를 그래프로 보여주며, 이 때 사용된 실제 데이터는 직선, 데이터를 통해 예측한 값은 점선으로 나타낸다.

2.3 진행 정도에 따른 예측모델 비교 분석

SIR 모형을 기반으로 구현한 예측모델을 사용하여, 초기 1달간의 데이터와 현재까지 진행된 데이터를 통해 예측을 실시한다. 예측 대상 나라는 확진자 표본이 많은 나라(미국, 러시아)와 방역 수준이 높다고 평가 되는 나라(이스라엘, 한국)들로 구성하였다. 현재(2020.5.29) 기준으로 미국 1,771,631명,

러시아 387,623명, 이스라엘 16,887명 한국 11,402명의 확진자가 발생하였다.

다음은 두가지 데이터 종류를 통해 예측한 그래프이다. 미국의 경우 데이터 범위가 너무 커, 그래프 확진자 범위를 1/10으로 제공하였다.



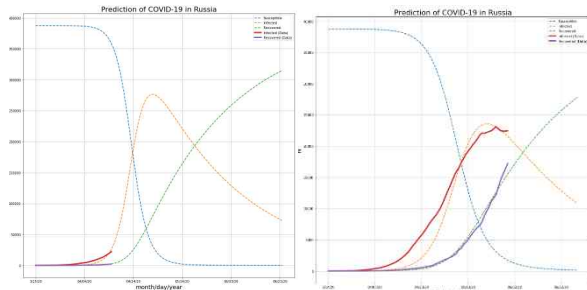
<Figure 2>Prediction of Confirmed / Recovered Cases in US
By initial(left) & Current Data(right)

<Table 1> Prediction of Active Cases in US Using Initial data(1 month) / Full data

Date	2/25	4/16	5/29	6/23
Prediction Using Initial Data	20	1,564,620	1,258,945	1,073,041
Prediction Using Full Data	20	317,500	1,148,270	865,488
Actual Active Cases	45	578,430	1,236,764	NaN

초기 데이터를 활용한 예측과 전체 데이터를 활용한 예측의 비교를 통해 확산이 초반에 급격히 이루어진 것을 확인 할 수 있다. 하지만 현재까지의 데이터를 통해 예측한 결과는 전체적으로 완만히 확산이 이루어 졌을 뿐 아니라 활성 확진자의 최댓값 또한 158 만 명에서 133 만 명으로 초기 예측 보다 낮아진 것(18%)을 확인 할 수 있다. 이와 같은 결과에서, 미국의 방역 시스템이 실행되지 않았던 초기에 비해, 현재까지 전염병 확산이 진행되면서 방역시스템이 확산 속도 및 확진자 수에 영향을 끼쳤다는 것을 알 수 있다.

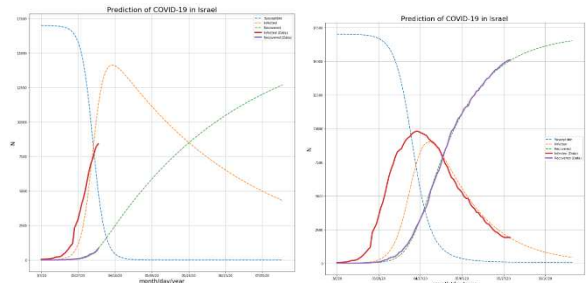
러시아와 이스라엘 데이터를 통해서도 해당 국가의 방역 시스템이 확산의 기율기 정도와 활성 확진자의 최댓값에 영향을 끼친다는 것을 알 수 있다. 러시아의 최댓값은 27 만명에서 23 만명으로 감소(17%)하였으며, 이스라엘의 최댓값은 1 만 4 천명에서 9 천명으로 감소(55%)하였다.



<Figure 3> Prediction of Confirmed / Recovered Cases in Russia
By initial(left) & Current Data(right)

<Table 2> Prediction of Active Cases in Russia
Using Initial data(1 month) / Full data

Date	3/15	4/23	5/29	6/23
Prediction Using Initial Data	20	164,667	147,828	74,358
Prediction Using Full Data	20	21,090	218,930	125,310
Actual Active Cases	55	57,320	223,990	NaN

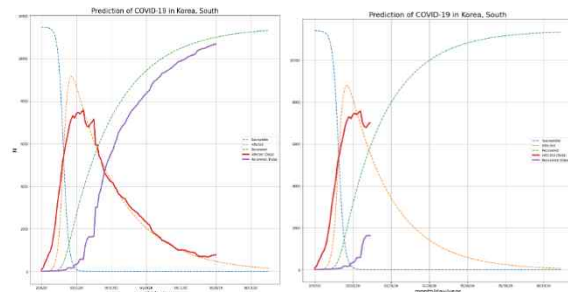


<Figure 4> Prediction of Confirmed / Recovered Cases in Israel
By initial(left) & Current Data(right)

<Table 3> Prediction of Active Cases in Israel
Using Initial data(1 month) / Full data

Date	3/8	4/27	5/29	6/23
Prediction Using Initial Data	2	12,533	8,189	5,863
Prediction Using Full Data	2	8187	1986	587
Actual Active Cases	55	8,151	1,927	NaN

한국의 경우 조금 특이한 형태의 그래프 형태를 가진다. 초기의 급격한 확산에 대응하여, 적극적인 방역 실시, 확산을 억제하는 것에 성공하였다. 그렇기 때문에 높은 수준의 방역시스템에도 불구하고 초기의 데이터와 현재까지의 데이터로 예측한 값들의 차이가 거의 나지 않는다.



<Figure 5> Prediction of Confirmed / Recovered Cases in South Korea
By initial(left) & Current Data(right)

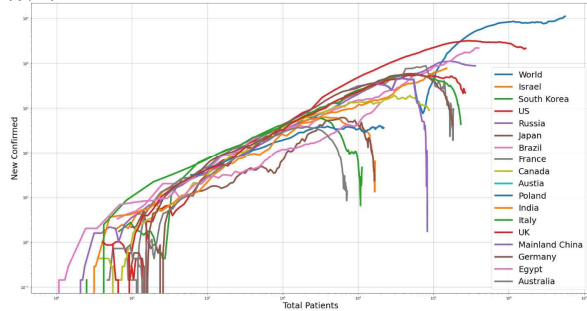
<Table 4> Prediction of Active Cases in South Korea Using Initial data(1 month) / Full data

Date	2/20	4/10	5/29	6/23
Prediction Using Initial Data	4	2,219	258	86
Prediction Using Full Data	4	2950	494	199
Actual Active Cases	87	3125	774	NaN

위와 같은 예측 비교를 통해 각 나라의 감염자 확산에 방역시스템이 영향을 끼쳤다는 것을 확인할 수 있었다. 그 중 이스라엘과 같이 방역시스템의 변화가 두드러져 시각적으로 확인할 수 있는 나라도 있는 반면, 한국이나, 미국처럼 한눈에 확인하기 어려운 나라 또한 존재한다. 그렇기 때문에 추가적인 방법을 통해 방역시스템의 수준을 확인한다.

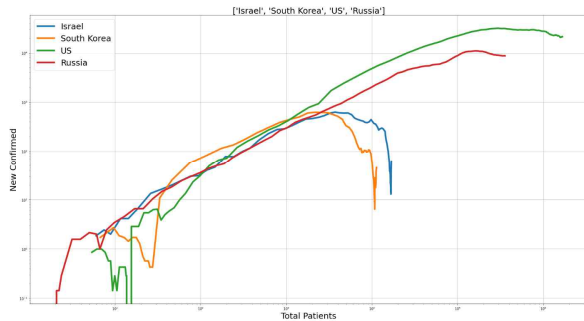
추가확진자와 누적 확진자 사이의 상관관계 그래프를 확인하면 각 나라 별 확진자 확산에

대한 기울기가 일정하다는 것을 확인 할 수 있다. 이것을 통해 코로나 바이러스의 확산 정도가 특정한 값을 가지고 있을 것이라고 추측 할 수 있다. 또한 기울기의 유지 기간을 통해 각 나라의 방역 시스템이 해당 확산을 얼마만큼 효과적으로 진압 하였는지 시각적으로 간략하게 확인 할 수 있다.



<Figure 6> Correlation Between Additional and Cumulative Confirmed

이전의 예측한 4 개 나라에 대한 예측의 오차 값과 해당 그래프에서 나타나는 기울기 지속 정도의 차이를 비교한다.



<Figure 7> Correlation Between Additional and Cumulative Confirmed in Four Country

해당 그래프를 통해 한국이 가장 효과적으로 방역 시스템을 실시하였고, 이스라엘, 러시아, 미국 순으로 기울기의 크기나 지속 정도가 커지고 있다는 것을 알 수 있다. 이와 같은 경향성이 이전의 예측 모델의 오차 값과 비슷한 것 또한 알 수 있다.

3. 결론

본 논문에서는 LSTM 과 SIR 모형 기반의 COVID - 19 확진자 예측 모델을 구현했다. 각 모델에서 구한 예측 값과 5/29 일의 실제 확진자 값을 비교하였다.

<Table 5> Accuracy of Prediction Using LSTM and SIR model

5/29 기준	Actual data	SIR	Error rate(%)	LSTM	Error rate(%)
Israel	16987	16928	0.3	16992	0.03
South Korea	11441	11477	0.3	11569	1.1
US	1746019	1717279.	1.65	1880718	7.7
Russia	387623	368527	4.9	393067	1.4

해당 표에서 SIR 모델과 LSTM 모델의 실제 확진자와 예측 값의 차이를 확인하였다. 또한, 해당 날짜의 실제 확진자와 예측 값의 오차 정도의 평균을 구하였을 때 정확도 측면에서 SIR 모델과 LSTM 모델의 결과값이 유의미한 차이를 보이지 않았다. 하지만, 사용되는 입력 값이나, 목적에 따라 적절한 모델 사용이 필요하다.

<Table 6> Average Accuracy of Prediction Using LSTM and SIR model

	SIR	LSTM
Israel	0.96%	1.12%
South Korea	1.75%	1.38%
US	4.35%	6.47%
Russia	4.94%	4.43%
Average	3.00%	3.35%

LSTM 모델의 경우, 다양한 input 값이 제공되어 있을 경우와 단기에 대한 예측에 강점을 가지고 있다. SIR 모델의 경우, 이전의 확진자, 회복자, 사망자에 관한 데이터만으로 예측이 가능하다. 대신 회복자 추세에 영향을 받아 정밀한 확진자 예측이 어려워 대략적 예측 값을 유추할 때 유리하다.

SIR 모델을 초기 데이터와 현재까지의 데이터를 구분하여 예측을 실시하였다. 두 예측 값의 차이를 통해 각 나라별 방역 수준의 정도를 예측해 본 결과, 이스라엘과 한국의 경우, 초기 예측과 현재 예측의 차이가 유의미하게 드러나는 것으로 보아 방역에 대한 여러 시도가 상당부분 성공적이었다는 것을 알 수 있다. 반면에, 러시아와 미국의 경우 앞의 차이 값에 비해 낮은 차이를 보이는 것으로 보아 방역이 비교적 전염병 확산에 비교적 적은 영향을 준 것으로 보인다.

참고문헌

GERS, Felix A.; SCHMIDHUBER, Jürgen; CUMMINS, Fred. Learning to forget: Continual prediction with LSTM. 1999.

Lee Mi Rim. 「」 Composition of a disease diffusion simulation model for the suppression of the spread of infectious diseases and the search for countermeasures 201612

NDIAYE, Babacar Mbaye; TENDENG, Lena; SECK, Diaraf. Analysis of the COVID-19 pandemic by SIR model and machine learning technics for forecasting. arXiv preprint arXiv:2004.01574, 2020.

D2.3 코로나19 접촉 추적 앱의 스트레스 요인이 그 유용성에 미치는 효과

Stressful Factors of Contact Tracing Applications Affecting Their Potential Capability to Mitigate the Spread of COVID-19

주재훈 (Jaehun Joo)

소속(국문): 동국대학교 상경대학 경영학부(정보경영학전공)

소속(영문): Department of Information Management, Dongguk University-Gyeongju

E-mail: givej@dongguk.ac.kr

한위밍 (Yumin Han)

소속(국문): 동국대학교 대학원 국제비즈니스협동과정

소속(영문): International Business Cooperative Course, Dongguk University

E-mail: 875816309@qq.com

Corresponding author: Jaehun Joo

123 Dongdae-ro, Seokjang-dong, Gyeongju 38084, Korea

Tel: +82-10-9299-7150, Fax: +82-54-770-2310, E-mail: givej@dongguk.ac.kr

국문 초록

코로나19의 대유행과 확산은 모든 사람들에게 큰 스트레스를 주고 있다. 디지털 접촉 추적 앱이 코로나19의 확산을 방지하는데 유용한 한편, 그 사용자들에게 스트레스를 유발하는 요인이 되기도 한다. 본 연구는 접촉 추적 앱의 유용성을 극대화하면서 사용자들의 스트레스를 줄이는 대처방안을 실증분석한다. 앱을 사용함으로써 유발되는 스트레스 요인, 앱 사용 과정에서의 평가, 사용자의 정서적 적응 행동, 그 잠재력을 발휘하여 사용하는 인퓨전의 관계를 구조방정식 모형으로 제안하였다. 연구모형에서 도출한 4개의 연구가설을 중국의 건강코드 사용자들을 대상으로 설문조사를 실시한 자료를 통해 연구가설을 검증하였다. 연구결과, 앱 사용자의 프라이버시 침해에 대한 염려를 제거함으로써 앱의 잠재력이 극대화되는 수준에 활용될 수 있다. 이는 결국 접촉 추적 앱이 코로나19를 종식시키는데 유용한 도구가 된다는 것을 의미한다.

주제어

COVID-19, 코로나19, Contact tracing app, Coping theory, Stress, Infusion

D2.4 의료 AI 시스템 개발을 위한 AI 윤리의 기술적 적용과 ‘설명가능한 AI(XAI)’ 구현 전략 —코비드-19 진단모델 중심—

안종훈
인공지능콘텐츠LAB
hamletahn@gmail.com

Abstract - 본 연구의 목적은 인공지능 기반 의료용 진단모델 개발 시 AI윤리의 기술적 적용 방법으로써 설명가능한 AI 기술을 어떻게 구현할 것인지 그 구현전략을 살펴보는 데 있다. 연구를 위해 코비드-19 진단에 사용되었던 연구모델들을 참고로 할 것이다. 연구진행을 위해 AI윤리의 적용 범위와 그 기술적 구현을 위한 ‘설명가능한 AI(XAI)’의 기술적 구현 요소를 먼저 살펴보고 그 요소들을 AI기반 의료용 진단모델 개발과정에 적용시킬 것이다.

Key Terms - 의료AI시스템, 진단모델, AI윤리, XAI, XAIP

I. 서론

인공지능(Artificial Intelligence) 기술이 국가기관과 기업운영을 말할 것도 없고, 국방 무기개발 및 의료산업 등 전 산업에 확산되고 있어, 관련 산업 주체자들은 AI 기반 시스템 개발에 모든 에너지를 집중하고 있다.

문제는 ‘어떻게’ 유용한 AI시스템을 개발할 것인가이다. 기술적으로 해당 산업과 연결된 AI시스템을 개발하는 것도 중요하지만 그 결과물은 해당 기관 뿐 아니라 사회적 국가적으로 해를 끼치지 않고 안전하고 책임있는 기술이어야 한다. 여기에서 AI윤리가 등장하였고, 그것을 기술적으로 적용시키는 방법으로써 ‘XAI(Explainable AI)’ 즉, 설명가능한 AI’(DARPA, 2017)가 본격적으로 등장하였다.

특히, 국방 무기개발 및 의료진단용 AI 모델과 기기들은 사람의 생명과 직결되는 분야로 그 윤리적 개발이 직접적으로 요구되어지고 있다. 2015년 5월 <Nature>에 게재된 ‘로봇공학: 인공지능의 윤리’는 치명적인 자동화 무기시스템인 ‘LAWS(lethal autonomous weapon systems)’를 우려하는 인공지능 선구자 스튜어트 러셀(Stuart Russell)의 주장을 담고 있다.

그런데, AI개발에 있어 기계학습과 딥러닝에서 학습데이터와 학습과정 그리고 평가 후 최종결과물이 나왔을 때 그 결과물의 구체적인 생성과정을 알 수 없어 ‘블랙박스’로 취급되어버리는 한계점을 노출시켰다. 예로써, 코비드-19 환자의 진단결과 90% 양성이라는 확진 결과가 나왔는데, 그것이 환자의 증상이나 영상촬영물을 종합하여 구체적 설명없이 기존

양성환자의 X-레이 등 영상물과 비슷하여 양성이라 최종진단을 내린다면 폐렴, 폐결핵 혹은 폐암 그리고 사스(SARS)나 메르스(MERS) 같은 다른 전염병의 가능성을 배제해 버리는 한계로 오진의 가능성이 발생하게 된다.

이런 한계를 극복하기 위해 설명가능 XAI와 XaiP기술을 개발하게 된 것이다. 이들 기술은 본질적으로 AI시스템의 설명가능성에 중점을 두고 있지만, 기술적 구현을 통해 투명성과 신뢰성, 안전성 그리고 책임성을 확보할 수 있어 추상적 AI윤리의 현실적 실천에 큰 기여를 하게 되었다.

II. AI윤리의 필요성과 기술적 적용

2020년 2월 28일 로마교황청에서는 AI 윤리 백서 ‘로마콜(Rome Call for AI Ethics)’을 발표하고 인공지능 시스템이 준수해야할 규정과 원칙을 발표했다. ‘로마 콜’은 AI가 윤리적으로 설계되어야 한다는 ‘윤리적 알고리즘(algo-ethical)’이라는 새로운 개념도 제안하였다. 윤리적 알고리즘을 사용하라는 의미이다.

특히 권리부분에서 AI의 윤리적 이용을 위해 필요한 6가지 원칙으로 투명성, 포용성, 책임성, 불편부당, 신뢰성, 보안과 프라이버시 등을 제시하였다.

국내의 경우 AI윤리의 필요성에 대해서는 이미 2017년에 기업들을 비롯한 우리 사회가 인공지능 윤리를 왜 중요하게 다뤄야 하는지에 대한 필요성을 설명하고, 지능정보사회 윤리 가이드라인(안)(김명주, 2017) 등 여러가지가 제시되었다. 당시 상황에서 가이드라인이지만 2020년 현재 상황에서는 가이드라인과 원칙을 넘어 구체적-기술적 방법이 제시되어야 할 것이다.

III. XAI의 구성요소 및 기술적 구현

XAI 구성요소

XAI는 그 개념상 최종 출력물 관련 설명에 관한 내용 뿐 아니라 여러가지 요소적 의미를 함축하고 있는 복합적 개념이다. 첫째는 학습과정과 결과에 대한 해석가능성과 설명가능성이고, 둘째는 시스템 개발 전체 과정에서의 투명성과 신뢰성이며, 세번째는 최종 결과물에 대한 책임성과 그에 따르는 안전성과 보안성 등의 개념이 모두 들어 있다.

또, 여기에는 기획자와 개발자, 프로그래머와 사용자 등 개발 생태계 속 모든 관계자들(stakeholders)이 포함되어 있다. 중요한 것은 이들 추상적 개념들을 AI 개발 과정 전체에 어떻게 기술적으로 구체화시킬 수 있는가에 있다.

XAI의 기술적 구현방안: XAIP 설계와 가버넌스

사실 AI 시스템을 개발하는 과정에는 데이터나 모델개발, 학습과정 설계 그리고 테스트 등 일련의 과정에는 기업이나 조직의 영업비밀들이 포함되어 있다. 따라서 기업의 영업비밀을 훼손시키지 않으면서 설명가능한 AI 시스템을 개발해야 하는 데, 이런 점을 보완하고 XAI의 목적 달성을 위해서는 학습모델 기반 XaiP(Explainable AI Planning) 설계의 개념(M. Cashmore et al, 2015)을 도입하고 있기도 하다. 특히 기업의 창의성과 영업비밀을 보호하기 위해서는 법적 제도적 관점에서 네거티브 규제와 감사 및 감독을 포함한 통합적인 가버넌스 시스템이 필요하다.

동시에 AI 시스템 개발의 밑단에 있는 컴퓨팅 기술 측면에서도 병렬과 분산 컴퓨팅 방법을 이용하여 AI윤리의 가이드라인과 윤리적 요소들을 중간중간에 하나의 노드(node)로써 융합시킬 수 있는 다양한 방법들을 고안해야 할 것이다.

XAI 구현을 위한 활용 가능한 툴(Tool)

산업에서 참고할 수 있는 XAI를 기술적으로 구현할 수 있는 기본적인 툴로는 다음과 같은 것들이 있다.

- IBM AIF 360
- Google XAI
- AI SimMachine
- H2O XAI
- MotionBoard for SkyAI
- Fiddler LABS AI

IV. 의료분야 XAI 기술적 구현사례: 코비드-19 진단모델 중심

사례 1: COVIDX-net: A Framework of Deep Learning Classifiers to Diagnose COVID-19 in X-Ray Images(March 24, 2020)

사례 2: MantisCOVID: Rapid X-Ray Chest Radiograph and Mortality Rate Evaluation with AI for COVID-19(May 4, 2020)

V. 결론 및 제언

AI 윤리 옹호자들은 ‘설명가능성’을 AI 시스템에 책임을 물을 수 있는 가장 중요한 방법 중 하나라고 한다. AI 시스템의 설명가능성 속에 개발과정의 투명성과 신뢰성, 최종 결과물의 안전성과 보안성 등 전반적인 AI 윤리가 포함될 수 있으므로 앞에서 언급한 XaiP설계와 학습모델의 철저한 구성을 통해 각 산업군의 AI개발에 참고하고 관련 전문인력 양성에도 관심을 기울여야 할 것이다.

VI. 참고문헌

김명주, “인공지능 윤리의 필요성과 국내외 동향”, 2017.

DARPA, “Explainable AI program”, 2017
Ezz ED Hemdan et al., “COVIDX-net: A Framework of Deep Learning Classifiers to Diagnose COVID-19 in X-Ray Images”, 2020.

Russel, Sturt, “Robot Engineering: the Ethics of Artificial Intelligence”, NATURE, May, 2015.

Maria Fox et al., “Explainable Planning”, Kings College London, 2017.

Michael Cashmore et al., “Towards

Explainable AI Planning as a Service”, Kings College London, 2015

Yagmur Yasar, “MantisCOVID: Rapid X-Ray Chest Radiograph and Mortality Rate Evaluation with AI for COVID-19”, 2020.

XAIP 2018, ICAPS Workshop on Explainable AI.
<http://icaps18.icaps-conference.org/xaip>
외

E2 [학술논문] 예측과 의사결정

E2.1 강화학습 주식투자 수익률 비교에 대한 연구

이현재
가톨릭대학교 경제학과
jaylee4274@gmail.com

이홍주
가톨릭대학교 경영학부
hongjoo@catholic.ac.kr

Key Terms: Reinforcement learning, Trading, Machine learning, Deep learning

I. 서론

주가는 해당 시점의 모든 정보가 반영되어 있다는 효율적 시장 가설에 따라 주가를 예측하는데 부정적인 시각이 많았다. 그러나 최근 기계학습 분야의 인공지능경망 기술이 발달하며 시장 평균 수익률을 초과할 수 있다는 연구 결과가 제시되고 있다.

본 연구는 기본 주가 데이터에 더해, 주가에 영향을 미친다는 거시경제 데이터를 이용하여 강화학습 주식투자 모듈에 적용하였을 때의 강화학습 방법간, 투자기법간 수익률을 비교한다. 주가 데이터는 KOSPI시장에서 몇 가지의 기준에 따라 8개의 그룹을 선정하여 구성하였고, 각 그룹엔 30개씩의 종목으로 구성하였다. [표1]

각 종목의 데이터는 2016.01.01년부터 2019.12.31년까지의 기본 주가 데이터와, 기타지표를 수집하였다. 거기에 추가적으로 거시경제 변수 4가지를 수집해 데이터셋을 구성하였다. 기본 주가 데이터(날짜, 시가, 고가, 저가, 종가, 거래량) 외의 기타지표와 거시경제 변수는 다음과 같다. [표2]

학습기간은 2017.01.01~2018.12.31로 2년간의 주가 데이터를 사용했다. 학습된 모델을 2019.01.01~2019.12.31의 기간에 테스트했다. 학습 데이터의 주가데이터는 전처리해 입력하였고, 대부분 5, 10, 20, 60, 120일의 이동평균 비율로 바꿔 단, 중장기의 추세를 학습에서 파악할 수 있도록 하였다.

강화학습 주식투자에서는 주어진 차트 데이터를 환경으로 한다. 에이전트는 주어진 환경에서 [매수, 매도, 관망] 3가지 행동 중 하나를 하게 되는데 이 실험에서는 100번의 반복학습을 하며 투자 행위에 대한 결과를 수익률로 산출해 이전의 투자 행위가 긍정적인지, 부정적인지 충분한 경험을 쌓도록 한다. LSTM 신경망을 사용해 학습하였다.

8개 그룹의 각 30개 종목에 대해 총 11가지 투자 기법을 적용해 수익률을 비교하였다. [표3]

강화학습 기법으로는 DQN(Deep-Q-Network), PG(Policy Gradient), AC(Actor-Critic), A2C(Advantage Actor-Critic), A3C(Asynchronous Advantage Actor-Critic)의 5가지 기법을 사용하였다. 기술적분석 지표 투자 기법으로는 SmaCross1, SmaCross2, SmaCross3, SmaCross4, RSI14를 사용하였다. SmaCross란 특정기간의 이동평균선이 교차하는 지점을 매매지점으로 잡는 투자기법을 말하며, SmaCross1은 10일, 20일의 이동평균선을 사용했고 SmaCross2는 20, 60일의 이동평균선을 사용했다. SmaCross3,4는 SmaCross1,2에 일정한 손절매 기준을 정한 기법을

의미한다. RSI14는 2주간의 시장 과열도에 대한 지표를 의미하는데 특정 값을 기준으로 매매지점으로 잡는 전략으로 사용했다. 마지막으로 5번의 무작위 투자의 평균값을 산출해 총 11개 방법간 수익률을 비교하였다. 그렇게 해서 강화학습 투자 기법의 가능성을 확인하고, 특정 종목 그룹에서 어떤 투자 기법이 더 높은 수익률을 보이는지, 그 차이가 통계적으로 유의한지 검증하는 것이 본 연구의 주제이다.

II. Figures and Tables

<표 1> 종목그룹 선정 기준

Index	Criteria	Group name
1	시가총액	대형주 그룹
2	시가총액	소형주 그룹
3	거래량	거래량 많은 그룹
4	거래량	거래량 적은 그룹
5	등락률	등락률 높은 그룹
6	등락률	등락률 낮은 그룹
7	등락률 절대값	등락률 절대값 높은 그룹
8	등락률 절대값	등락률 절대값 낮은 그룹

<표 2> 추가 변수

Index	Additional index	Description
1	GNP	외국인순매수량
2	FR	외국인보유비율
3	INP	기관순매수량
4	KOSPI	코스피 지수
5	S&P	S&P 500 지수
6	WTI	텍사스 중질유 현물가격
7	USD	원달러환율

<표 3> 사용한 투자 기법

Index	Trading strategy	Description
1	DQN	DQN 방식의 강화학습 투자
2	PG	PG 방식의 강화학습 투자
3	AC	AC 방식의 강화학습 투자
4	A2C	A2C 방식의 강화학습 투자
5	A3C	A3C 방식의 강화학습 투자
6	SmaCross1	10, 20일 이동평균선 전략
7	SmaCross2	20, 60일 이동평균선 전략
8	SmaCross3	SmaCross1 + 4ATR 손절 전략
9	SmaCross4	SmaCross2 + 4ATR 손절 전략
10	RSI14	30이하 매도, 70이상 매수

11	MTrader	무작위 투자
----	---------	--------

III. 참고문헌

Park , S. C. “The Moving Average Ratio and Momentum”, The Financial Review 45(2010), 415-445.

김선웅, 최홍식, “외국인 거래정보를 이용한 트레이딩시스템의 성과분석”, 경영과학 32(4), 2014.12, 57-67

김문권, 파이썬과 케라스를 이용한 딥러닝/강화학습 주식투자(도서)

E2.2 LSTM분류의 KOSPI 지수 예측에의 적용

이재철
서울과학기술대학교
산업정보시스템학과
zhqhddl13@naver.com

천세학
서울과학기술대학교
경영학과
shchun@seoultech.ac.kr

Abstract - 인공지능망, SVM 등 머신러닝을 이용한 주가예측에 관한 연구는 지난 20년동안 많이 진행되어 왔다. 최근 시계열 데이터를 다루는데 LSTM모델이 많이 주목받고 있다. 본 논문에서 LSTM을 적용하여 KOSPI 지수를 분류한다. 분류한 결과를 cut-off 지점의 변화를 통해 수익률이 어떻게 달라지는 지 살펴본다. 또한 지수 분류 시 추가적인 정보량이 어떻게 영향을 주는지 단순모델, 경제변수모델, 기술적 지표모델의 수익률을 비교 분석한다.

I. 서론

- 머신러닝을 이용한 주가예측에 관한 연구는 지난 20년동안 많이 진행되어왔다.
- 인공지능망, SVM, k-Nearest Neighbor 등 다양한 방법론을 주가예측에 적용해왔다.
- 특히 최근 딥러닝을 통한 주가예측이 주목받게 되었고, 주가와 같은 시계열 데이터의 경우 과거 시계열 데이터와 관련성을 염두에 둔 LSTM분류 모델이 주목 받고 있다.(Do-Hyung Kwon et al.,2019)
- 본 논문에서는 LSTM을 적용하여 주가를 분류하고자 한다.
- 특히 이진분류 모델에서 중요하게 사용되는 정확도, 정밀도, 재현율, 특이도 등의 지표를 사용하여 cut-off 지점의 변화를 통해 수익률이 어떻게 변하는지를 알아보고자 한다.
- 또한 어떤 cut-off 지점이 공격적인 투자를 하는지 보수적인 투자를 결정하는지 여러 지표를 통해 알아보고자 한다.
- 주가가 하락할 것이라고 예측했는데 실제 주가가 상승하여 수익의 기회를 놓친 것보다는 주가가 상승할 것이라고 예측했는데 실제 주가가 하락하여 손실을 보게 되는 경우 일반적으로 투자자입장에서 더 회피하고 싶다.
- 따라서, 본 논문에서 LSTM을 통해 분류할 때, 이러한 False Positive 에러를 줄이면서 과감한 투자 전략을 세우는 법을 여러 방법을 통해 분석하고자 한다.
- 또한 지수 예측 시 추가적인 정보량이 어떻게 영향을 주는지 단순모델, 경제변수모델, 경제변수와 기술적 지표모델의 수익률을 비교 분석한다.

II. 기존연구

기존 머신러닝을 통한 주가 예측에 관한 연구는 다음과 같다.

2.1 기존 문헌 연구

정보 및 변수 사용 측면

- 1) Bao et al. (2017): 시가, 종가, 고가, 저가, 거래량과 기술적 지표를 변수로 이용
- 2) 변태우 외 (2011): 환율과 같은 거시경제지표를 변수로 사용.

2.2 Confusion matrix, ROC curve

Confusion matrix는 분류의 성능을 평가하는 행렬이다. Confusion matrix의 행은 실제 클래스, 열은 예측한 클래스를 나타냅니다. 이진분류 문제의 경우 각 4개의 케이스를 TN(예측과 실제 모두 0),FP(1이라 예측했지만 실제 0),FN(0이라 예측했지만 실제 1),TP(예측과 실제 모두1)로 나누어 평가한다.

ROC curve(Receiver Operating Characteristic): 분류 모델의 성능을 평가하기 위한 그래프로, 기준 값(0과 1을 나누는 기준값인 cut-off)에 따른 성능을 나타낸다. FPR을 x축으로, TPR을 y축으로 하여 그려진 그래프를 수치적으로 표현하기 위해 그래프 아래 면적을 AUC(Area Under the Curve)라 하는데 이 면적은 클수록 좋다.

FN과 FP의 경우 낮을수록 좋은 모델이지만 많은 경우에 있어서 error cost가 다른 경우들이 많다. 의료문제에 있어서 암환자를 진단할 때 암환자라고 예측했는데 아닌 것보다, 암환자가 아니라 했는데, 암환자인 경우 환자에게 더 큰 피해를 끼칠 수 있다.

금융 거래 문제에 있어서 이상 거래라고 예측했지만 아닌 경우보다 이상 거래가 아니라고 예측했지만 이상 거래였던 경우가 회사엔 막대한 피해를 끼칠 수 있다.

III. Cut-off Model 과 실험의 소개

3.1 모델 소개

기존 연구들과 차이점 (단순 accuracy만 높아서는 수익률이 항상 높다고 얘기할 수 없다.)

Fall-out은 FPR(False Positive Rate)

특이도(Specificity) = $TN / N = (1 - FP / N)$

3.2 실험데이터 소개

2005년 1월부터 2020년 5월 20일까지의 코스피 지수의 시계열 데이터

- 변수의 소개:

기본적 코스피 지수 정보: 코스피의 시가, 종가, 상한가, 하한가, 거래량

거래량 세분화: 개인, 외국인, 기관 투자자별 누적 매매량

기술적 지표: 이동평균선, Macd, Dmi, 볼린저 밴드
거시경제지표: 환율(달러, 유로, 엔, 위안화), 금, 원유, 구리, 회사채3년 금리

- 예측 모델의 종류

모델1=(기본적 코스피 지수 정보)

모델2=(모델1 + 거래량 세분화)

모델3=(모델2 + 기술적 지표)

모델4=(모델2 + 거시 경제 지표)

모델5=(모델2 + 기술적 지표 + 거시경제 지표)

기간 1 = 1일의 수익률이 양수일 때 1, 음수일 때 0

기간 3 = 3일의 수익률이 양수일 때 1, 음수일 때 0

기간 5 = 5일의 수익률이 양수일 때 1, 음수일 때 0
 기간 7 = 7일의 수익률이 양수일 때 1, 음수일 때 0
 기간 10 = 10일의 수익률이 양수일 때 1, 음수일 때 0
 기간 14 = 14일의 수익률이 양수일 때 1, 음수일 때 0

- 예측 변수 및 분석의 내용 설명:
 1 단계) LSTM 모델을 사용하여 총 3765개의 데이터 중 2700개를 train dataset, 700개를 validation dataset, 마지막 365개의 데이터를 test dataset으로 분류, 출력층을 sigmoid 함수로 지정하여 0과 1로 바꾸어 종가의 방향 분류
 2 단계) 훈련한 모델을 test dataset에 적용시키고 Accuracy를 최대화하는 cut-off 값, Recall - Specificity를 최소화하는 cut-off 값, 기본 Default로 0.5로 지정된 값 세 개의 모델로 분류
 3 단계) 각 결과값과 실제 값을 confusion matrix를 이용하여 성능을 평가

IV. 실험 결과 논의

4.1 실험결과 분석

이번 이진분류 모델에서 줄여야 할 목표는 False Positive(1이라 예측했지만 실제로는 0)이다.

정밀도(Precision) = $TP / (TP + FP)$

특이도(Specificity) = $TN / (TN + FP)$

재현율(Recall) = $TP / (TP + FN)$

Precision과 Specificity가 둘 다 높을수록 좋지만 극단적인 경우 Precision과 Specificity는 높지만 성능은 좋지 않은 경우 발생. 따라서 Recall도 고려를 해야 한다.

(부연 설명) 특이도(Specificity)는 $TN / (FP + TN)$ 이고, 위양성율(Fall out), 즉 FPR(False Positive Rate)은 $FP / (FP + TN)$ 이기에 1-위양성율이 바로 특이도이다. Specificity와 Recall은 trade-off 관계에 있기 때문에 두 가지 모두 고려하여 모델을 선택한다.

ACC: Accuracy 최대화, RS: Recall-Specificity 최소화, Def: cut-off=0.5

<표1> AUC

	1일	3일	5일	7일	10일	14일
모델1	0.59	0.62	0.66	0.74	0.78	0.79
모델2	0.58	0.62	0.66	0.73	0.76	0.78
모델3	0.54	0.60	0.62	0.66	0.64	0.58
모델4	0.58	0.62	0.64	0.69	0.71	0.70
모델5	0.56	0.58	0.64	0.72	0.70	0.70

<표2>ACC

단 위 %	1일	3일	5일	7일	10일	14일
모델1	31.43	31.10	30.58	30.39	26.81	23.89
모델2	13.31	26.43	30.66	22.70	26.19	19.52
모델3	3.39	26.18	23.43	19.14	23.59	4.11
모델4	20.54	20.30	26.65	24.06	12.26	11.51
모델5	21.91	31.57	17.45	14.50	11.06	24.82

<표3>RS

단 위 %	1일	3일	5일	7일	10일	14일
모델1	26.11	28.57	26.72	36.96	27.16	30.73
모델2	28.00	21.48	30.29	24.61	26.31	29.89
모델3	38.63	28.65	14.82	17.34	16.69	20.19
모델4	28.51	19.43	22.86	10.25	18.14	14.77
모델5	20.46	40.76	25.01	13.34	20.56	20.89

모델1	29.49	29.52	31.20	29.40	25.17	27.53
모델2	27.49	23.56	31.28	25.75	26.41	22.06
모델3	33.12	18.27	16.10	13.50	14.88	19.66
모델4	16.57	20.15	25.71	18.80	10.18	14.63
모델5	29.61	36.19	20.00	16.41	10.19	22.99

<표4>Def

단 위 %	1일	3일	5일	7일	10일	14일
모델1	26.11	28.57	26.72	36.96	27.16	30.73
모델2	28.00	21.48	30.29	24.61	26.31	29.89
모델3	38.63	28.65	14.82	17.34	16.69	20.19
모델4	28.51	19.43	22.86	10.25	18.14	14.77
모델5	20.46	40.76	25.01	13.34	20.56	20.89

<표5>Def의 Precision

	1일	3일	5일	7일	10일	14일
모델1	0.6538	0.7500	0.7463	0.8727	0.9615	0.8980
모델2	0.5874	0.6242	0.6777	0.8644	0.7667	0.8235
모델3	0.5906	0.6723	0.6424	0.6748	0.6667	0.5867
모델4	0.6429	0.6279	0.6474	0.7570	0.8028	1.0000
모델5	0.6053	0.5745	0.7222	0.8384	0.8537	0.8140

<표6>Def의 Specificity

	1일	3일	5일	7일	10일	14일
모델1	0.8902	0.9157	0.8976	0.9573	0.9879	0.9689
모델2	0.4817	0.6446	0.7651	0.9512	0.8303	0.8882
모델3	0.6280	0.7651	0.6446	0.7561	0.7333	0.6149
모델4	0.7256	0.8072	0.6325	0.8405	0.9152	1.0000
모델5	0.6341	0.5181	0.8193	0.9018	0.9273	0.9012

<표7>Def의 Recall

	1일	3일	5일	7일	10일	14일
모델1	0.1700	0.2143	0.2577	0.2474	0.2632	0.2316
모델2	0.6050	0.5000	0.4227	0.2629	0.4842	0.4421
모델3	0.4400	0.4082	0.5464	0.4278	0.4632	0.4632
모델4	0.4050	0.2755	0.5773	0.4154	0.3000	0.1799
모델5	0.4600	0.5510	0.4021	0.4256	0.3684	0.3704

4.2 결과 논의 및 시사점

- 모델 3,4의 AUC도 낮고 수익률도 저조하다.
- 기간 수익률을 3일과 5일로 잡았을 때 대체로 수익률이 높다.
- 가장 기본 모델인 Def도 괜찮은 성능을 보여준다.

V. 결론

- Def 모델은 꾸준한 수익률을 보여준다. 하지만 정밀도와 특이도는 높지만 재현율은 낮다. 이는 실제 1 중에 1로 예측한 비중이 낮다는 것을 의미한다.
- 따라서Def 모델은 보수적인 투자자가 선택하기에 알맞은 전략이다. 하지만 1년 동안 선택한 횟수가 적기 때문에 AUC가 높은 모델도 동시에 고려하는 것이 바람직하다.
- AUC가 높다고 항상 수익률이 높은 것은 아니다. 그러나 각 모델 별, 기간 별 수익률의 변동성은 낮아진다.

VI. 참고문헌

Bao W, Yue J, Rao Y (2017) A deep learning

framework for financial time series using stacked autoencoders and long-short term memory. PLoS ONE 12(7): e0180944

변태우, 이성우, 김재광, 이지형 (2011). 텍스트마이닝과 환율데이터를 활용한 주가의 위험성 예측에 관한 연구. 한국지능시스템학회

Do-Hyung Kwon, Ju-Bong Kim, Ju-sung Heo, Chan-Myung Kim, and Youn-Hee Han(2019) Time Series Classification of Cryptocurrency Price Trend Based on a Recurrent LSTM Neural Network , Journal of Information Processing Systems Vol. 15, No. 3, pp. 694-706, Jun. 2019

E2.3 Factors influencing Investment Decision of Individual Investors in Social Impact Investment using Crowdfunding platform

Pham Duong Thuy Vy
순천향대학교 경영대학
yunjikim1219@sch.ac.kr

Tran Hung Chuong
순천향대학교 경영대학
hyunjoon193@sch.ac.kr

최재원
순천향대학교 경영대학
jaewonchoi@sch.ac.kr

Abstract – Recently, social impact investing is becoming a hot issue in the social enterprise’s field. Investors around the world are using impact investments to utilize radically the power of capital. Growing impact investment market provides capital to address the world’s most pressing challenges in sectors such as climate change, sustainable agriculture, renewable energy, conservation, etc and affordable and accessible basic services including housing, healthcare, and education. At the same time, the core of the financing plan for social enterprises has shifted from the government-based “support system” to the “investment-type” by the private sector. Besides that, crowdfunding is also becoming a funding solution for entrepreneurs. Using crowdfunding, entrepreneurs are able to call for capital from individual investors which is easier than calling for capital from big companies or organizations. Therefore, this study will consider crowdfunding as a potential platform in the social investment market. By that, we want to clarify the role of crowdfunding in social enterprises and do individual investors willing to contribute capital for impact investment social enterprises if crowdfunding being used as a platform? And we also want to find out which factors influencing investment decision of individual investors in social impact investment using crowdfunding platform. Furthermore, to complete this study, survey and interview will be used as main research methodology. The purpose of conducting a survey that is utilized when a specific target involved; interview method is to explore the responses of the people to gather more and deeper information that is related to this study. We plan to conduct a survey of 500 entrepreneurs of small and medium-sized businesses who are tend to invest in social impact field to clarify which factors influence their investment decision. This study is expected to confirm the role of crowdfunding in social impact enterprises; to make it easier for entrepreneurs in social impact investment field to access capital from individual investors through crowdfunding platform in the future.

Key Terms: Crowdfunding, Impact investment, Impact economy, Individual investor, Social economy, Social impact.

E2.4 재무 정보 기반의 기계 학습을 활용한 투자 의사결정

전소영
한양대학교
경영대학 경영학부
thdud1282@hanyang.ac.kr

김동성
한양대학교
경영대학 경영학부
paulus82@hanyang.ac.kr

전상경
한양대학교
경영대학 경영학부
sjun@hanyang.ac.kr

김종우
한양대학교
경영대학 경영학부
kjlw@hanyang.ac.kr

Abstract - 본 연구에서는 안정적인 수익 확보를 위한 중장기적 투자를 지원하고자 재무 변수를 투자 지표로 활용하였다. 이를 위하여, 기업 재무 데이터를 수집하여, 학습 모델별 예측 성능을 비교한 후 주가수익률 상승의 클래스에 대한 확률이 높은 종목들을 대상으로 주식 거래 시뮬레이션을 수행하였다. 재무 정보 기반의 기계 학습을 통해 실질 수익률을 얻을 수 있다는 분석 결과를 바탕으로 중장기 투자 지표의 가능성을 발견할 수 있다.

Key Terms - 기계 학습, 재무 정보, 투자 의사결정, Neural network, Prediction of probability

I. 서론

오늘날 주식 투자자들은 불확실성이 높은 금융 시장에서 수익 창출을 목적으로 기본적 분석과 기술적 분석에서 나아가 기계 학습 기법의 활용 방안에 대해서 연구하고 있다. 특히, 기계 학습을 활용한 연구들은 주로 단기간의 주가 예측을 중심으로 이뤄지고 있다(Park et al., 2016). 그러나 기업 외부의 정치·사회적 이슈로 급변하는 금융 시장에서 단기 투자는 안정적인 수익률을 보장하지 않는 전략이 될 수 있다. 이에 본

예측하는 이진 분류기로써 기계 학습을 적용하였다.

입력 값으로 활용한 재무 데이터는 총 26개로 기존 연구(Kim et al., 2008; Ballings et al., 2015)에서 활용된 6개 변수와 신규 재무 변수 20개 변수(연율화한 영업 이익의 Trendline, 총자산회전률의 품질, ROIC의 안정성 등급 등)로 나뉠 수 있다. 본 연구에서는 기존 연구에서 활용된 재무 변수의 데이터 셋(# of features =6)과 신규 재무 변수를 포함한 데이터 셋(# of features =26)을 나눠 학습을 진행했다.

III. 분석 방안 및 결과

본 연구에서 재무 정보를 활용한 분석 방안은 다음과 같다. 우선 1650개 기업을 대상으로 2013년 3월부터 2018년 9월까지 공시된 재무 정보와 주식 데이터를 수집한 후, 기업별 26개의 변수와 3개월/6개월 간 주가수익률을 측정하였다. 다음으로 2013년 3월부터 2017년 6월까지의 데이터를 훈련하였고, 이후 기간의 데이터에 대해 예측정확도를 확인하였다(<표 1> 참고).

<표 1> 예측 결과 - 기간별 예측 정확도(%).

Period	Declining rate of return	# of features	Random Forest	Gradient Boosting	Support Vector Machine	Neural Network	Logistic Regression
3개월	54.3%	6	0.526	0.530	0.533	0.535	0.540
		26	0.535	0.541	0.531	0.540	0.536
6개월	52.4%	6	0.515	0.521	0.523	0.519	0.525
		26	0.524	0.517	0.519	0.514	0.523

코스피·코스닥 시장 내 1650개 종목을 대상으로 2013년 3월부터 2018년 9월까지 공시된 재무 정보와 3개월 및 6개월 간 주가수익률을 수집하였으며, 기계 학습을 활용한 시뮬레이션을 통해 실질 수익률을 확인하였다. 이를 바탕으로, 안정적인 투자를 위한 지표 개발의 가능성을 발견할 수 있다.

II. 기계 학습 및 데이터 설명

기계 학습은 컴퓨터를 학습시켜 분류나 예측에 활용하는 기술이다. 본 연구에서는 총 5가지 모델(랜덤 포레스트(Random Forest), 그래디언트 부스팅(Gradient Boosting)과 서포트 벡터 머신(Support Vector Machine), 로지스틱 회귀분석(Logistic Regression), 인공 신경망(Artificial Neural Networks))의 기계 학습을 진행하였다. 재무 정보를 입력 값으로 사용하여 주가수익률의 방향(상승 또는 하락)을

<표 2> 시뮬레이션 - 상위 종목의 평균 수익률(%).

추가로, 기계 학습을 통해 각 클래스에 대한 확률 값을 예측한 결과를 바탕으로 주식 거래 시뮬레이션을 수행하였다. 2017년 9월부터 2018년 6월까지 1년의 기간 중에 3개월, 6개월 주기로 주가수익률 상승의 확률이 높은 상위 3개, 5개, 10개 종목을 해당 주기마다 매수, 매입을 진행한 후 수익률을 비교하였다(<표 2> 참고).

IV. 결론

본 연구에서는 기계 학습을 통해 기업의 재무 정보를 기반으로 주가수익률의 등락을 예측하는 것이 가능한지 확인하였다. 그러나 재무 정보를 활용한 기계 학습은 모든 모델에서 해당 주기의

수익률 하락(%)보다 낮은 예측 정확도를 보였다는 점에서 유의미하지 않은 예측력을 보였다. 따라서, 지속적인 수익을 확보하기 위한 주가 예측의 취지를 반영하여 예측력이 확실한 기업만 선정하여 수익률을 높이는 전략으로 투자 시뮬레이션을 진행하였다. 분석 결과, 코스피·코스닥 지수보다 높은 수익률을 보였으며 특히 인공 신경망을 활용했을 때 더욱 우수한 결과를 확인할 수 있었다. 또한 20개의 신규 재무 변수가 포함된 데이터 셋일 때, 그리고 3개월보다 6개월을 주기로 할 때 더 높은 수익률이 산출되었다. 이를 바탕으로, 주식 투자자의 목적을 달성하기 위한 의사결정을 지원하기 위한 투자 지표의 가능성을 발견할 수 있다. 향후 주가 예측의 한계를 보완하는 재무 변수를 개발하여 6개월 이상의 장기 투자를 위한 연구를 진행할 필요가 있다.

Period	# of features	Random Forest	Gradient Boosting	Support Vector Machine	Neural Network	Logistic Regression
3개월	6	2.82%	13.66%	4.73%	-1.60%	1.34%
	26	3.38%	15.87%	4.77%	27.11%	8.87%
6개월	6	4.03%	27.87%	7.54%	2.47%	4.69%
	26	15.54%	16.93%	9.36%	56.05%	2.58%

참고문헌

Ballings, M., D. Van den Poel, N. Hespeels, and R. Gryp, "Evaluating multiple classifiers for stock price direction prediction," *Expert Systems with Applications*, Vol.42, No.20(2015), 7046-7056.

Heo, J. Y., and J. Y. Yang, "SVM based Stock Price Forecasting Using Financial Statements," *KIISE Transactions on Computing Practices*, Vol.21, No.3 (2015), 167-172.

Kim, K. Y. and K. R. Lee, "A Study on the Prediction of Stock Price Using Artificial Intelligence System," *Korean Journal of Business Administration*, Vol. 21, No.6(2008), 2421-2449.

Kim, S. D., "Data Mining Tool for Stock Investors' Decision Support," *The Journal of the Korea Contents Association*, Vol.12, No.2(2012), 472 - 482

Park, J. Y., J. P. Ryu, and H. J. Shin, "Predicting KOSPI Stock Index using Machine Learning Algorithms with Technical Indicators," *Journal of Information Technology and Architecture*, Vol.13, No.2(2016), 331-340.

Tsai, C.F., Y. C. Lin, D.C. Yen, and Y. M. Chen, "Predicting stock returns by classifier ensembles," *Applied Soft Computing*, Vol.11, No.2(2011), 2452-2459.

A3 [기획세션] 디지털 비즈니스 II

A3.1 RPA 적용 사례 및 AI

이우항 - KSTEC 전문위원

A3.2 인공지능으로 인한 금융비즈니스 탈바꿈

김엄상 - 투이컨설팅 상무

B3 [특별세션] 인텔리전스 대상 기업세션 II

B3.1 머신 러닝 기반의 자동차 파손 및 심도 인식 모델

김태운 - (주)애자일소다 AI알고리즘연구팀장

B3.2 하이퍼오토메이션 실현을 위한 인공지능 기술과 AI InspectorOne

김계관 - (주)그리드원 대표이사

B3.3 Redesigning Digital Banking with AI

강정석 - (주)에이젠글로벌 대표이사

B3.4 집에서 하는 셀프구강검진 '이아포'

양희영 - (주)큐티티 선임연구원

B3.5 인공지능 QA 자동화 서비스 HBsmith

한종원 - (주)에이치비스미스 대표이사

B3.6 노트의 Vision AI 모델 자동 압축 기술

이기환 - (주)노타 이사

C3 [특별세션] 인공지능 응용 경진대회 II

C3.1 SOM-K모델을 이용한 과목별 학습전략 제시(자기조절학습 변인을 중심으로)

원종관, 장영진, 김수민, 김민수, 올가 체르냐예바(부산대)

C3.2 K-Means와 XGBoost를 사용한 이커닝 플랫폼 이용자 이탈방지 전략 연구

김태영, 조예린, 이세한, 김소연, 윤소정, 이세한(연세대)

D3 [학술논문] 개인화 기술

D3.1 추천의 다양성을 고려한 그래프 기반 추천시스템

나혜연

UNIST

융합경영대학원

hyna@unist.ac.kr

남기환

한국과학기술원 경영대학

namkh@kaist.ac.kr

Abstract - 추천시스템 서비스는 사용자의 성향에 맞는 추천을 제공할 때 만족도를 높일 수 있다. 본 연구는 추천시스템의 다양성 확보가 사용자의 취향에 따라 다르게 받아들일 수 있음을 가정하여 다양성 추구자를 구분하고 사용자 특성에 맞게 그래프 기반의 추천시스템을 적용하고자 한다.

Key Terms - Recommender system, Diversity, Graph-based network.

I. 서론

추천시스템의 주요 목적은 인터넷 환경에서 고객이 편하고 빠르게 구매를 할 수 있도록 유도하여 기업에게 이익을 높이며 동시에 고객은 구매를 통한 만족감을 얻게 하는 것이다. 실제로, 추천시스템의 이러한 효과는 추천시스템을 가진 온라인 쇼핑몰의 성장을 통해 증명되었다. 최상위 e-commerce 책임자들은 추천시스템이 고객의 선택에 매우 중요한 영향력을 가진다고 언급했는데, 그 예로 넷플릭스에서 미디어 선택의 60%, 아마존 매출의 약 35%가 추천시스템에서 나왔다는 사실이다 (Dokyun Lee and Kartik Hosanagar 2019). 다양한 종류의 추천기법이 있지만, 협력필터링은 추천시스템 중에서 현재까지 가장 우수한 성능을 나타낸다고 알려진 기법으로써, Goldberg et al.(1992)에 의해 그 개념이 처음으로 소개되었다 (Jieun Son et al., 2015). 협력필터링은 사용자와 비슷하게 상품을 보거나 구매하는 다른 유저를 찾거나 혹은 비슷하게 보거나 구매한 유저의 상품을 추천하는 방식이다 (Schafer JB et al., 1999). 하지만, 협력필터링을 기반한 추천시스템이 기업과 고객을 지속적으로 만족시키기 힘든 부분도 있음을 확인하고 있다. (Fleder D and Hosanagar K, 2009)에서 분석 모델과 시뮬레이션을 통해 협력필터링이 평균적으로 많이 구매하는 제품을 집중적으로 추천하는 문제가 있음을 발견했다.

또한, 충분한 평점 확보가 되지 않는 경우 역시 협력필터링이 한 쪽으로 치우친 추천의 결과가 나온다는 문제도 지적되고 있다 (Mooney and Roy, 2000). 불완전한 Collaborative Filtering 이 거대한 하나의 취향을 만들고 사람들을 그쪽으로 유인한다고 볼 수 있는 지점이다. 기업 측면에서 장기간으로 보면, 롱테일 현상이 일어나게 되고, 기업의 매출은 감소하게 된다. 고객 측면에서도 고객의 취향이 고정되어 있지 않기 때문에 만족스러운 추천이 되지 않는다. 고객의 취향이 변화하는 것은 다양한 측면에서 파악할 수 있는데, 일정 부분은 사용자의 특

성이고 또 다른 부분은 환경적 요인 때문이다 (Yehuda Koren, 2009). 예를 들어, 동일한 영화를 감상해도 몇 년 전에는 낮은 평점을 줬으나 최근에는 결혼, 자녀, 직업 변화 등 변화된 상황으로 인해 높은 평점을 줄 수 있다. 추천시스템이 사용자에게 만족을 주기 위해서는 사용자의 아이템 선택 성향을 보다 잘 파악할 수 있어야 한다. 사용자의 취향은 갑자기 변화하기도 하고 혹은 사용자 스스로 취향을 확신하지 못해 지속적으로 넓은 폭의 아이템 쇼핑을 즐길 수도 있다. 기존에 오프라인 상품 마케팅을 위해 사용자의 구매 성향을 파악하고자 하는 연구가 있었다 (cobb and Hoyer, 1986). 하지만, 추천시스템을 사용하는 사용자의 구매 성향에 대한 연구는 활발하지 못했다. 본 연구에서는 추천시스템이 기업과 고객의 만족도를 더 높이기 위해, 고객의 성향을 찾은 뒤 그에 맞게 제품을 추천하는 것에 초점을 두었다. 특별히 고객의 성향을 시간의 흐름에 따라 파악하여 변화하는 고객의 성향을 추천시스템에서 파악할 수 있도록 계획하였다. 고객의 성향은 취향의 다양성을 의미하며, 다양성 확보를 위해 다양성 추구자를 구분한다. 다양성 추구자는 단기간, 장기간, 채널을 나누어 확인하였으며, 이러한 사용자들은 그래프 기반의 추천시스템을 통해 유저의 성향에 맞는 파라미터 최적화 작업이 진행된다. 그래프 기반의 추천시스템은 여러 방식의 추천시스템 모형을 반영시키며 발전시킬 수 있고 (Zan Huang et al., 2002), 다양한 성격을 반영시킬 수 있는 유연성이 있어 (Lei Shi, 2013) 여러 연구에서 활용되고 있다.

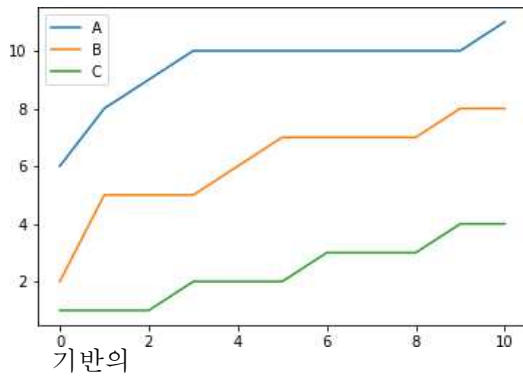
I. 연구방법

분석에는 국내 기업 패션 추천시스템의 클릭기반 데이터를 사용하였다. 클릭기반 데이터를 사용한 이유는 구매 기반 데이터에 비해 사용자의 이동 경로를 상세하게 파악할 수 있다는 점에서 사용자의 구매 성향을 정확하게 이해할 수 있기 때문이다.

서론에서 언급한 바와 같이, 본 연구에서는 모든 사용자들 사이에서 다양성 추구자를 구분해내는 과정이 중요하다. 사용자들은 내적 혹은 외부적인 동기로 인해 다양성 추구자가 될 수 있는데, 이러한 변화는 시간의 흐름과 병행한다. 따라서, 개인의 다양성 추구 인덱스를 시간의 흐름과 결합하여 파악하되, 시간을 단기와 장기로 나누어 다양성 추구의 지속성을 점검하였다.

<그림 1> 사용자들의 단기 다양성 차이

분류된 다양성 추구자는 그래프



추천시스템에 적용된다. 그래프 기반의 추천시스템은 유저 네트워크 생성과 아이템 네트워크 생성 후, 네트워크 투영을 통해 유저에게 아이템을 추천하게 된다. 사용자 네트워크에서는 다양성 추구 인덱스로 노드 연결의 강도를 달리하며, 아이템 네트워크는 제품별 품목의 차이에 따라 강도를 달리한다.

본 연구에서는 시간의 흐름에 따른 사용자의 다양성 추구 변화, 다양성에 맞는 네트워크 구성 기법으로 추천시스템의 다양성을 얻고자 하는 것으로 요약할 수 있다.

I. 참고문헌

Dokyun Lee, Kartik Hosanagar, How Do Recommender Systems Affect Sales Diversity? A Cross-Category Investigation via Randomized Field Experiment, *Information Systems Research*, Vol. 30, No. 1, March 2019, pp. 239–259.

Goldberg, K., Roeder, T., Gupta, D., and Perkins, C. (2001), Eigentaste: A constant time collaborative filtering algorithm, *Information Retrieval*, 4(2), 133–151.

Jieun Son, Seoung Bum Kim, Hyunjoong Kim, Sungzoon Cho, Review and Analysis of Recommender Systems, *Journal of the Korean Institute of Industrial Engineers*, Vol. 41, No. 2, pp. 185–208, 2015.

Schafer JB, Konstan J, Riedl J (1999) Recommender systems in ecommerce. *Proc. 1st ACM Conf. Electronic Commerce (ACM, New York)*, 158–166.

Fleder D, Hosanagar K (2009), Blockbuster culture's next rise or fall: The impact of recommender systems on sales diversity. *Management Sci.* 55(5):697–712.

Mooney RJ, Roy L (2000) Content-based book recommending using learning for text categorization. *Proc. 5th ACM Conf. on Digital Libraries (ACM, New York)*, 195–204.

Yehuda Koren (2009), Collaborative Filtering with Temporal Dynamics, *KDD '09: Proceedings of the 15th ACM SIGKDD international conference on Knowledge discovery and data mining*, Pages 447–456.

Cathy J. Cobb, Wayne D. Hoyer, Planned versus impulse purchase behavior, *Journal of retailing*, Vol. 62.1986, 4, p. 384–409.

Zan Huang, Wingyan Chung, Thian-Huat Ong, Hsinchun Chen, A Graph-based Recommender System for Digital Library 2002. 65–73, *Proceedings of the Second ACM/IEEE-CS Joint Conference on Digital Libraries*.

Lei shi, Trading-off Among Accuracy, Similarity, Diversity, and Long-tail: A Graph-based Recommendation Approach, 2013, pp. 57–64, *Proceedings of the 7th ACM conference on Recommender systems*.

D3.2 AI 대 전문가의 추천 효율성 비교 실험 연구: 이커머스 내 패션 스타일 추천을 중심으로

김주영

한국과학기술원 경영대학

gloriakim@ kaist.ac.kr

남기환

한국과학기술원 경영대학

namkh@kaist.ac.kr

Abstract - 급변하는 기술환경에서 인공지능(AI)은 제조업, 서비스업, 소매업, 마케팅 등 다양한 분야에서 비즈니스를 변화시키고 있다. 이와 같은 AI의 발전은 기계가 사람에게 도전하는 프로젝트의 제기를 제공했으며 이에 따라 인간의 일자리 대체에 대한 위협을 일으켰다. 이에 대해 창의성 또는 공감 능력과 같은 독특한 인간의 속성을 필요로 하는 분야에서는 인간의 일자리가 늘어날 것이라는 설득력 있는 견해가 제기되었다. 하지만, 최근 AI 연구가 창의성(Creativity)과 공감지능(Empathetic Intelligence) 등의 영역까지 확대되면서 창의적인 업무와 서비스업까지도 AI 대체 가능성 논쟁이 일고 있다. 기존 연구들이 이커머스 실험연구를 통해 대체제 추천의 효과나 개선된 알고리즘으로 인한 성능 차이를 다루었다면, 본 논문은 전통적으로 창의적 업무(creative task)로 인지도, 패션 스타일링 추천 서비스를 AI가 설계한 것인지, 인간 전문가가 설계한 것인지에 따라 소비자의 반응이 어떻게 달라지는지 연구한다. 이는 AI와 사람에 대한 사용자의 인지 효과를 컨트롤 하기 위해 추천 주제를 공개하지 않고 이뤄졌으며, 이를 위해 3,120,934 명의 참가자를 대상으로 대규모 랜덤화 현장 실험을 실시했다. 이는 많은 이커머스 사이트들이 상대적으로 소홀했던 보완제 기반의 스타일 추천 서비스의 필요성 조망하고, 인간 전문가들이 제공해온 창조적인 업무를 인공지능 알고리즘 사용하여 새롭게 접근한다. 이를 통해 AI 대 인간 전문가의 경쟁적인 측면을 조망하고 중요한 경영적 및 학문적 의미를 산출하고자 한다.

Key Terms - 인공지능, 추천시스템, 이커머스, 온라인 패션 비즈니스, 패션 산업, 패션 스타일링, 기계적 창의성, 창의적 업무

I. 서론

AI가 적용되는 산업 및 업무의 범위가 늘어나면서 기계가 사람에게 도전하는 프로젝트에 대한 사람들의 관심이 증대되었다. 러시아의 체스 거장 Garry Kasparov가 1997년 IBM의 "Deep Blue" 컴퓨터에 패한 AI 자동화 혁명의 첫 번째 희생자가 된 이후, 연구진이 AI 기술 개발에 노력을 기울인 결과 AI는 여러 분야에서 인간 이상의 퍼포먼스를 보여주었다. 이에 더하여 언어 번역, 자율 주행, 챗봇을 활용한 고객 서비스 등에서 AI 적용 연구가 활발히 진행되고 있다. 이렇듯 인공지능 대 인간 지능에 대한 호기심이 증가하면서, AI 자동화로 인해 인간의 일자리가 위태로워지는 것 아니냐는 우려가 나오고 있다. 이에 대해 인간의 창의성이나, 공감 능력과 같은 인간의 속성을 필요로 하는 예술, 서비스 분야에서는 AI 자동화가 어려울 것이라는 설득력

있는 견해가 제기되었다. 하지만 최근 Computational Creativity와 공감 지능(Empathetic Intelligence)까지 연구범위가 확대되면서 회화, 작곡, 패션 디자인 등 창의적인 예술 분야와 AI 대체 가능성에 대한 논쟁이 일고 있다. 따라서 본 연구는 AI 기술이 가장 보편적으로 적용되고 있는 분야인 이커머스 추천 시스템 중, 창의적 업무(creative task)에 해당하는 패션 스타일링 추천 서비스에 대하여 AI 대 전문가의 퍼포먼스 비교를 시도한다. 창의적 업무인 패션 스타일링 추천 서비스를 AI가 설계한 것인지, 인간 전문가가 설계한 것인지에 따라 소비자의 반응이 어떻게 달라지는지 실제 실험을 통해 분석하는 것이다. 이는 AI와 사람에 대한 사용자의 인지 효과를 컨트롤 하기 위해 추천 주제를 공개하지 않고 이뤄졌으며, 이를 위해 3,120,934 명의 참가자를 대상으로 대규모 랜덤화 현장 실험을 실시하고 양측간 효율성을 객관적인 지표에 따라 비교 분석하였다. 따라서, "스타일 추천에 의한 제품 만족, 고객 선호 예측과 만족, 그리고 이커머스의 매출 향상"이라는 목적을 위해 인간 전문가들이 수행한 창조적인 업무를 AI 알고리즘을 사용하여 새롭게 접근하였다. 이는 AI 대 인간 전문가의 경쟁적인 측면을 조망했다는 점에서 중요한 경영적 및 학문적 의미를 산출한다.

III. 연구 설계

AI와 전문가의 효율성을 비교한다는 연구 목적에 맞게 사용자 선호도 및 매출&판매 성과를 비교하고자 하였다. 본 논문의 실험의 AI 스타일 추천은 Item2Vec 머신러닝 알고리즘이 적용되었다. 한편, 실험에 참여한 50 명의 패션 전문가는 Au et al(2004)에 의해 발전된 Analytical Framework of Fashion Design Process에 따라, 창의적인 업무를 수행했다고 보았다.

IV. 결과

<표 1> 결과 변수에 대한 결과 및 모델 성능 비교

Outcome Variable	AI	EP	AI -EPL	t-stat.	p-value
CTR	0.0139	0.00954	0.00437	5.154***	0.0000
CR	0.0031	0.00665	-0.00346	-3.963***	0.0002
SA	30189994	2759556	2594425	1.709*	0.0974

Note. *** $p < 0.01$, ** $p < 0.05$, * $p < 0.1$.

본 실험은 AI와 Expert의 퍼포먼스 비교라는 점에서 CTR(click through rate), CR(Conversion Rate), ASA(Average sales amount)라는 3가지 측정 지표를 검토하였다. 요약된 결과의 전체 개요는 <표 1>과 같다. CTR은 전체 노출 대비 클릭을 보여주는 지표로, AI와 Expert의 퍼포먼스는 매우 유의한

차이를 보였다. 이는 AI 가 Expert 보다 더 많은 클릭 수를 유도했다는 것을 의미한다. CR 퍼포먼스는 클릭당 구매 건수 전환을 보는 지표로 AI 와 Expert 의 추천간 유의미한 차이를 보였으며, AI 의 CR 은 Expert 가 AI 보다 더 많은 클릭당 구매 전환을 일으켰다. 총 판매액(Total Sale Amount)의 지표에서는 AI 가 더 많은 수익을 이끌어 낸 결과가 유의미 하였다. 즉, 이 결과를 통해 우리는 AI 가 expert 에 비해 User 들에게 더 많은 클릭(CTR)을 유발 시키지만, 클릭을 통한 구매 건 전환(CR)은 적게 유발 시킨다. 뿐만 아니라 이커머스 비즈니스의 수익에 있어서 가장 중요한 지표인 총 판매액(SA)은 AI 가 Expert 를 능가한다는 것을 사실을 도출하였다.

III. 참고문헌

Acemoglu, D., & Restrepo, P. (2018). Artificial intelligence, automation and work (No. w24196). National Bureau of Economic Research.

Adomavicius, G., & Tuzhilin, A. (2005). Toward the next generation of recommender systems: A survey of the state-of-the-art and possible extensions. *IEEE Transactions on Knowledge & Data Engineering*, (6), 734-749.

Agrawal, A., Gans, J. S., & Goldfarb, A. (2019). Artificial intelligence: the ambiguous labor market impact of automating prediction. *Journal of Economic Perspectives*, 33(2), 31-50.

Autor, D. H., & Dorn, D. (2013). How technology wrecks the middle class. *The New York Times*, 24(2013), 1279-1333.

Campbell, M., Hoane Jr, A. J., & Hsu, F. H. (2002). Deep blue. *Artificial intelligence*, 134(1-2), 57-83.

Grace, K., Salvatier, J., Dafoe, A., Zhang, B., & Evans, O. (2018). When will AI exceed human performance? Evidence from AI experts. *Journal of Artificial Intelligence Research*, 62, 729-754.

Guan, C., Qin, S., Ling, W., & Ding, G. (2016). Apparel recommendation system evolution: an empirical review. *International Journal of Clothing Science and Technology*. Huang, M. H., & Rust, R. T. (2018). Artificial intelligence in service. *Journal of Service Research*, 21(2), 155-172.

Luce, L. (2019). Generative Models as Fashion Designers. In *Artificial Intelligence for Fashion* (pp. 125- 139). Apress, Berkeley, CA.

Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Frontiers: Machines vs. humans: The impact of artificial intelligence chatbot disclosure on customer purchases. *Marketing Science*, 38(6), 937-947.

Silver, D., Huang, A., Maddison, C. J., Guez, A., Sifre, L., Van Den Driessche, G., ... & Dieleman, S. (2016). Mastering the game of Go with deep neural networks and tree search. *nature*, 529(7587), 484.

Silver, D., Schrittwieser, J., Simonyan, K., Antonoglou, I., Huang, A., Guez, A., ... & Chen, Y. (2017). Mastering the game of go without human knowledge. *Nature*, 550(7676), 354-359.

Vermeulen, B., Kesselhut, J., Pyka, A., & Saviotti, P. (2018). The impact of automation on employment: just the usual structural change?. *Sustainability*, 10(5), 1661.

Intelligence and Information Systems, Vol.11,

No.3(2005), 69~81.

Breiman, L., "Heuristics of Instability in Model Selection," Technical Report, Statistics Department, University of California at Berkeley, 1994.

Kim, C. S. and J. H. Lee, *Syntactic Pattern Recognition*, Prentice-Hall, New Jersey, 1985.

Au, J. S., Taylor, G., & Newton, E. W. (2004). Model of design process of Hong Kong fashion designers. *Journal of Textile and Apparel Technology and Management*, 4(2), 1- 4.

E3 [학술논문] 콘텐츠 분석

E3.1 댓글 분석을 통해 제한된 콘텐츠와 제한되지 않은 콘텐츠를 분류---Youtube으로 중심

장수평
순천향대학교 경영대학
zhangxiuping@sch.ac.kr

여방
순천향대학교 경영대학
luvfang6@sch.ac.kr

최재원
순천향대학교 경영대학
jaewonchoi@sch.ac.kr

Abstract - 세계적인 최대의 비디오 검색 및 공유 플랫폼인 Youtu는 사용자가 비디오를 업로드하거나 공유할 수 있을 뿐만 아니라 비디오에 댓글을 달 수 있다. 청소년을 보호하기 위해 YouTube는 연령 제한 콘텐츠에 대한 정책을 정의했다. 그러나 연령 제한으로 표시되지 않은 제한된 콘텐츠가 여전히 많이 있다. 청소년은 성인보다 편견과 유해한 콘텐츠에 부정적인 영향을 받을 가능성이 높기 때문에 청소년의 온라인 안전을 보호하기 위해 제한된 콘텐츠와 제한되지 않은 콘텐츠를 분류하는 것이 중요하다. Youtube시청자는 일반적으로 자신의 감정이나 의견을 표현하기 위해 콘텐츠에 댓글을 달고 있다. 그리고 댓글은 콘텐츠 품질을 이해하는 중요한 요소 중 하나 일 수 있다. 따라서 본 논문에서는 콘텐츠 댓글을 분석하여 제한된 콘텐츠와 제한되지 않은 콘텐츠를 분류하는 방법을 제안한다. YouTube에서는 제한된 콘텐츠의 댓글과 제한되지 않은 콘텐츠의 댓글을 포함하여 두 개 데이터 세트를 수집하고 클리닝했다. Word2vec를 이용해서 댓글을 벡터로 표시하고 분류기를 설정한다. 본 연구는 제한된 콘텐츠와 제한되지 않은 콘텐츠를 분류하고 비디오 공유 사이트의 환경을 정화하며 청소년의 신체적 및 정신적 건강을 보호하는 데 긍정적인 영향을 줄 수 있다.

Key Terms: 콘텐츠 분류, Word2vec, 감정분석

E3.2 악성 댓글 분류 시스템 모니터링 연구: 네이버 클린봇 분석

유지웅
경희대학교 이과대학
wldnd4011@khu.ac.kr

황보유정
경희대학교 이과대학
hwangbo@khu.ac.kr

손동성
경희대학교 이과대학
ehdtjd92@khu.ac.kr

이경전
경희대학교 경영대학
klee@khu.ac.kr

Abstract - 인공지능 기기의 오작동 및 편향성 문제 등 여러 부정적 영향에 대해 우려하는 목소리가 적지 않다. 본 논문에서는 네이버의 악성 댓글 분류기 ‘클린봇’을 분석하여 안정적으로 악성 댓글을 분류하는지 확인한다. 네이버 뉴스에서 댓글과 답글을 포함한 90,541건의 댓글을 수집하였고, 악성 댓글 분류 시스템 모니터링 절차에 따라 분석하여 안정성을 확인하였다.

수집된 댓글 중 클린봇에 의해 차단된 댓글은 총 1802건이다. 하지만 연구원들이 직접 분류한 결과 False Positive는 2864건으로 클린봇이 882건의 악성 댓글을 차단하지 못한 것을 알 수 있었다. 또한 False Negative는 8건이 잘못 차단되었다. 즉, 클린봇은 댓글을 대체로 과소하게 차단한다고 추측할 수 있다. Verification단계에서는 클린봇이 어떤 모델을 사용하고 있는지 추론하였으며, 같은 댓글임에도 불구하고 다른 결과가 나타나는 것을 통해 확정적인 규칙 기반 모델을 사용하는 것은 옳은 것으로 판단하였다. 모델의 안정성을 평가하기 위해 댓글에서 사용된 비속어를 수집하여 비속어 사전을 구축하고, 이를 기반으로 변형된 비속어를 클린봇모델이 안정적으로 차단했는지 확인한 결과 변형된 비속어를 포함하는 악성 댓글의 차단율은 약 26.83%로 나타났다.

Key Terms - 공정성, 인공지능, 인공지능 모니터링

사사

이 논문은 2017년 대한민국 교육부와 한국연구재단의 지원을 받아 수행된 연구임 (NRF-2017S1A5B8059804)

I. 연구내용

인공지능 시스템이 공정한 의사결정을 내리는 문제가 점점 중요해지고 있다(김윤정, 2018). 본 연구는 네이버의 악성 댓글 분류 시스템인 클린봇이 공정하고 안정적으로 작동하고 있는지 악성 댓글 분류 시스템 모니터링 3단계를 통해 분석한다. 첫째, 네이버 뉴스의 댓글과 답글을 수집하고 전처리 기법을 통해 불필요한 데이터를 제거한다. 둘째, 사람의 가치판단 단계는 전체 댓글을 사람의 눈으로 직접 확인하여, 클린봇이 악성 댓글을 공정하게 분류하고 있는지 여부를 확인한다. 셋째, 수집된 댓글을 분석하여 클린봇 모델의 일관성 및 안정성을 확인한다. 위 과정을 통해 모델의 동작 방식을 추론하였다. 한글 비속어의 특징 중 하나인 변형 비속어에서도 안정적으로 동작하는지 검증하였다. 본 연구에서는 네이버 클린봇의 일관성 및 안정성에 대해 중요 쟁점으로 다루었으며, 편향성에 대해서는 향후 연구로 진행 예정이다.

i. 데이터 수집 및 전처리

본 연구팀은 네이버 클린봇 모델의 품질을 확인하기 위해 2019년 11월 12일부터 2019년 12월

14일까지 ‘문제인’과 ‘김건모’를 키워드로 총 360건의 뉴스 기사를 수집했다. 댓글과 답글 및 클린봇 차단 여부 등 네이버 클린봇 분석에 필요한 정보들을 수집했고, 네이버 뉴스 댓글 정책에 기존 욕설, 비속어 등 일부 단어에 대해서는 ‘000’로 표시된다. 위 경우 클린봇은 무조건 차단하는 것으로 확인되어 본 연구에서는 제외하였다. 그 외 작성자가 삭제한 댓글, 운영규정 미준수로 인해 삭제된 댓글 등 연구 분석에 적절하지 않은 댓글들을 제외한 총 90,541개의 댓글이 수집되었다.

ii. 사람의 가치판단(Validity)

데이터 수집 및 전처리 단계에서 수집된 90,541개의 댓글의 기반으로 클린봇의 품질을 분석하였다. 악성 댓글이지만 차단하지 못한 ‘False Positive’와 정상 댓글이지만 차단한 ‘False Negative’를 통해 Validity를 확인하였다. 본 연구팀은 약 9만개의 댓글을 눈으로 확인하여 댓글에 비속어 포함 유무에 따라 검증을 진행하였다. 클린봇이 차단한 악성 댓글은 총 1,802개로 전체 댓글 중 약 1.99%의 비율로 악성 댓글이 사용된 것으로 보이지만, False Positive를 통해 확인한 결과 악성 댓글의 개수는 2,684개로 약 2.96%의 비율로 악성 댓글이 사용된 것으로 확인되었다. <표 1>은 False Positive를 통해 도출된 결과의 예시로 이 댓글들은 악성 댓글이 아닌 것으로 분류되었지만, ‘쉐키’, ‘ㅂㅅ’ 등의 비속어를 사용하여 차단 되었어야 할 댓글들이다.

<표 1> False Positive 예

댓글
ㅂ처먹어라 공산당 매국노 쉐키야...
그럼너는 건모 따까리냐!!! ㅂㅅ

False Negative를 통해 확인된 댓글은 총 8건으로 대체로 창문이면서 내용이 부정적인 댓글이 많이 포함되었다. <표 2>는 False Negative의 예시 댓글이다. ‘조국같은님...’에서 ‘님’이 비속어인가의 판단이 다를 수 있으나, ‘정봉주 같은님’, ‘조국기 같은님’의 댓글에서는 차단이 되지 않은 것으로 확인되었다. 즉, <표 2>의 사례들은 단어 ‘조국’이 차단에 영향을 미쳤을 수도 있다고 판단되나, 비속어가 사용되지 않음에도 불구하고 차단된 사례로 판단하였다.

<표 2> False Negative 예

댓글
황씨야 당신이 조국 정도만 되면 자한당지지율 50% 넘는다.
조국같은님...

ii. 모델 분석(Verification)

본 단계에서는 모델의 동작 방식을 유추하고 모델이 한글의 특징 중 하나인 변형 단어에서도 안정적으로 동작하는지 확인한다. <표 3>은 같은 내용의 댓글이지만 서로 다른 결과를 보이는 예이다. 네이버는 2020년 3월 19일부터 댓글 이력을 전면 공개하여 특정 인물의 과거 댓글 이력을 확인 할 수 있다. 두 댓글은 유사한 내용의 뉴스 기사, 같은 사용자, 같은 내용의 댓글이지만 서로 다른 결과를 보이는 한 사례이다. 클린봇 모델이 확정적 규칙 기반으로 악성 댓글을 차단하는 모델일 경우 동일한 댓글은 항상 같은 결과로 도출될 것이다. 하지만 같은 댓글에 결과가 상이한 것을 볼 때, 확정적 규칙 모델이 아닌, 여러 다른 변수가 개입하는 딥러닝 기반의 분류 모델이 사용되었을 것으로 추정된다.

<표 3> 같은 댓글 다른 결과의 댓글 예

	적용 전	클린봇 적용 후
서울신문	lmia**** > 2019.12.10. 18:50 · 신고 위례기	lmia**** > 2019.12.10. 18:50 ① 클린봇이 부적절한 표현을 감지한 댓글입니다.
연합뉴스	lmia**** > 2019.12.09. 15:03 · 신고 위례기	lmia**** > 2019.12.09. 15:03 · 신고 위례기

다음으로는 한글의 변형 단어에 대한 모델의 안정성을 분석하였다. 한글의 경우 여러 가지 유형의 자모 변형 및 음절 변형, 축약 등이 포함되어 있다(Kang, 2014). 그러므로 여러 유형의 변형에도 클린봇 모델이 안정적으로 악성 댓글을 차단하는지 확인할 필요가 있는데, 수집된 뉴스 댓글에서 사용된 비속어를 수집하여 비속어 사전을 구축하였다. <표 4>는 수집된 댓글에서 가장 많이 사용된 비속어 및 변형 단어의 예이다. 구축된 비속어 사전을 기반으로 각 비속어가 다음과 같이 변형되었을 경우 클린봇 모델이 얼마나 안정적으로 댓글을 차단하는지 수집된 댓글을 기반으로 확인하였다.

<표 4> 비속어 사전 및 차단율

비속어	변형 단어	차단율
씨발	시발, 씨방, C8, 시바	60.49%
지랄	지럴, 지랄, 지 랄, 지라	7.36%
병신	등신, ㅂㅅ, ㅂㅈ, 빙신	39.71%
미친	미친, ㅁㅈ, 미틴	22.98%
쓰레기	씨레기, 쓰래기, ㅅㄹㄱ	11.48%
새끼	새키, 새퀴, 시키, 색히	24.25%

<표 4>의 비속어 및 변형 단어의 차단율을 분석하였다. 90,541개의 댓글에서 10번 이하의 회귀 사례의 변형 단어는 제외하였는데, 각 비속어의 차단율은 다음의 식으로 계산하였다.

$$\beta = \frac{\sum_{i=1}^N \frac{b_t}{b_t + b_f}}{N} \times 100$$

은 전체 댓글에서 10번 이상 나타나는 변형 단어의 개수로 비속어 및 변형 단어의 차단율은 클린봇에 의해 차단된 댓글 중 변형 단어가 포함된 댓글의 수와 정상 댓글에서 변형 단어가 포함된 댓글의 수의 비율 총합의 평균으로 연산된다. 이를 통해 각 단어 별로 차단율을

산출할 수 있었으며 네이버 클린봇은 약 26.83%의 차단율을 보였다.

II. 참고문헌

Kang, Seung-Shik. "A Normalization Method of Distorted Korean SMS Sentences for Spam Message Filtering." *KIPS Transactions on Software and Data Engineering* 3.7 (2014): 271-276.

김윤정, *인공지능 기술 발전에 따른 이슈 및 대응 방안*, 한국과학기술기획평가원, 2018.
Available at <https://www.kistep.re.kr/c3/sub3.jsp>
(Downloaded 15 may 2020).

E3.3 A Study on the User Perception in ASMR Marketing Content through Social Media Text-Mining: ASMRtist vs Normal Influencer

Tran Hung Chuong
순천향대학교 경영대학
hyunjoon193@sch.ac.kr

Pham Duong Thuy Vy
순천향대학교 경영대학
yunjikim1219@sch.ac.kr

최재원
순천향대학교 경영대학
hyunjoon193@sch.ac.kr

Abstract – Nowadays, Autonomous Sensory Meridian Response (ASMR) - the experience of tingling sensations in the crown of the head - is rapidly growing in popularity and increasingly appearing in marketing. Not event in TV commercial advertisement, ASMR also fast growing in one-person media communication, many brands and social media influencers use ASMR for their marketing contents. The purpose of this study is to measure consumers' perceptions about the introduced products in ASMR marketing content and compare the differences in advertising effects of ASMR content creator and normal content creator in the same Marco influencer tier - the YouTuber that has 10,000 – 1,000,000 subscribers. The research methods selected ASMRtist (YouTuber who use ASMR in their contents) and normal YouTuber, we collect text comments from 400 videos of tech-device review videos (ASMR contents - 100/normal contents - 100) and beauty-fashion videos (ASMR contents - 100/normal contents - 100). 55,600 texts from ASMR contents and 50,700 texts from normal contents were analyzed by applying the LDA topic modeling algorithm, content analysis. Through the result, we can know that ASMR contents take more viewers' attention because in the same influencer tiers ASMR contents are getting more views, likes and comments than normal contents. As the result of LDA analysis, we found that ASMR content viewers and normal content viewers have different perceptions about content and product. In the Tech-device reviews field, when normal contents viewers mainly focus on product, all the topics are talking about product's design, release, price, purchase , etc and the ASMR viewers also focus on the product like product's performance, purchase but there are many topics are talking about reaction of ASMR sound, trigger, relaxation. In the Beauty-fashion field, normal content viewers mostly focus on style, makeup, fashion, product topics and ASMR content viewers' topics mainly focus to the reaction of ASMR trigger, response to ASMRtist and other topics are talking about makeup - fashion, product, purchase. Via LDA analysis, un-like normal contents, comments are focus on the main fields, many ASMR viewers comment that they feel more comfortable when watching the marketing content that used ASMR. This result has shown that ASMR marketing contents have a good performance in term of user watching experience than the normal content, so apply ASMR will take more consumer intention and may behave better advertising acceptance. As an influencer marketing strategy, this study provides information to establish an efficient advertising strategy by using influencer that creates ASMR content.

Key Terms: ASMR marketing, influencer, social media, text mining, topic modeling

E3.4 Dynamic Topic Model을 활용한 미래 시점 특허 문서 생성에 관한 연구

양낙영
연세대학교
산업공학과
yny0506@yonsei.ac.kr

김영
연세대학교
산업공학과
0kimy@yons ei.ac.kr

양재원
연세대학교
산업공학과
jwyanq1993@gmail.com

김동욱
연세대학교
산업공학과
tmdh78@naver.com

임우담
연세대학교
산업공학과
woodam@yonsei.ac.kr

이태욱
연세대학교
산업공학과
dlxodnr22@yonsei.ac.kr

김우주
연세대학교
산업공학과
dlxodnr22@yonsei.ac.kr

Abstract – With the recent rapid development of science and technology, technological innovation through research and development has exploded across a wide range of fields. In this situation, in order to understand the rapid development of science and technology and keep pace with it, it is necessary to have material to clearly understand technological innovation, which is the patent document. Patent documents are materials that can clearly describe innovation activities in virtually all fields and have characteristic that it is accumulated over time in a long time series. If you understand patent documents that have these characteristics, you can understand the growth and decline of specific technologies. It is also possible to predict which technologies will emerge in the future. Accordingly, in this study, time-series information was extracted from a set of patent documents collected using a dynamic topic model, which is an unsupervised learning based model suitable for extracting time-series information. In addition, future time information was derived by applying this time series information to a time-series predictive model, and by using this, patent documents of future time points could be generated.

Finally, in order to verify whether the model clearly creates a patent document at a future time, we assumed a specific year as the future time point and generated patent documents at that time point. Then, we derived similar documents by comparing the similarity between the generated document and the patent documents that were applied at that time. As a result, it was possible to reliably verify the performance of this model, as documents with high similarities to the generated documents were clearly drawn.

Key Terms – Patent Analysis, Time Series Analysis, Dynamic Topic Model, Document Generation

I. 서론

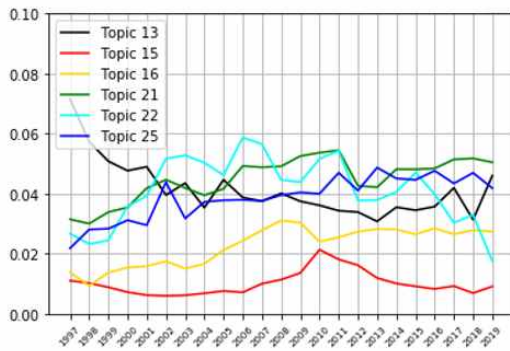
최근 과학 기술이 유례없는 급격한 발전을 맞이함에 따라 연구개발을 통한 기술적 혁신은 광범위한 분야에 걸쳐 폭발적으로 이루어지고 있다. 실제로 정부는 2020년도 국내 연구개발 예산을 전년도 대비 17.3% 증액된 24.1조원으로 편성하였는데, 이는 10년만의 두 자릿수 비율 증액으로써 국내 과학 기술 연구개발의 성장세와 그에 대한 두터운 기대를 한 눈에 보여준다. 이러한 상황에서 과학 기술의 급격한 발전을 이해하고 그에 발맞추어 기술적 진보를 이룩하기 위해서는 기술혁신을 명확히 파악할 수 있는 자료가 필요한데, 그에 해당하는 것이 바로 특허 문서이다. Griliches(1998)는 특허 문서가

실질적으로 모든 분야의 혁신활동을 명확히 설명할 수 있는 자료이며 시계열적으로 장기간 축적되어 있는 특성을 지닌다고 하였다. 즉, 특허 문서는 다양한 기술의 시계열적 변동에 대한 정보를 담고 있으며, 이러한 특허 문서를 이해한다면 특정 기술의 성장과 쇠퇴를 파악할 수 있고, 향후 어떠한 기술이 새롭게 출현할 것인지에 대한 예측 또한 가능할 것으로 판단된다(Kim and Bae, 2017).

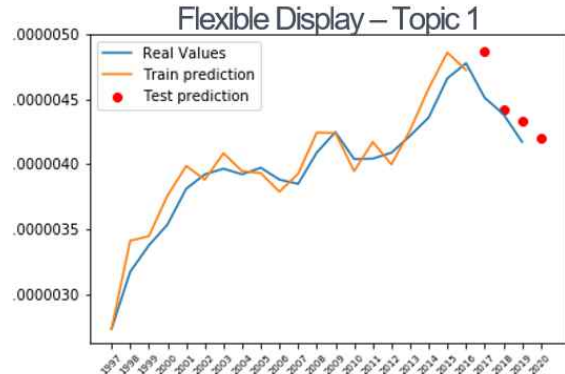
이에 따라, 본 연구에서는 시계열적 정보를 파악하기에 적합한 비지도 학습 기반 모형인 Dynamic Topic Model(Blei and John, 2006)을 활용함으로써 수집한 특허 문서 집합에서 다중 시점의 토픽 분포 및 토픽-단어 분포를 도출하였다. 이 때, 다중 시점의 토픽 분포 및 토픽-단어 분포는 시계열적 변동성을 지니는 시계열 분석 데이터로 판단할 수 있는데, 따라서 이러한 시계열 데이터를 시계열적 예측 모형에 적용함으로써 미래 시점의 토픽 분포 및 토픽-단어 분포 또한 도출할 수 있었다. 또한, 앞서 언급한 과정으로 예측한 분포는 Latent Dirichlet Allocation(Blei et al., 2003) 모형에서 가정하는 문서 생성 과정에 의거하여 새로운 문서 생성을 위한 입력 값으로써 활용될 수 있는데, 본 연구에서는 해당 방식을 통해 다양한 미래 시점의 문서를 생성하였다. 마지막으로, 본 모형이 미래 시점의 특허 문서를 명확히 생성하는 지를 검증하기 위해 특정 연도를 미래 시점으로 가정하였고 해당 시점의 특허 문서를 생성하였다. 그 뒤, 그 시점에 실제 출원되었던 특허 문서와 생성된 문서 간의 유사도를 비교하여 유사 문서를 도출하는 방식을 활용하였는데, 생성된 문서와 높은 유사도를 보이는 문서들이 명확히 도출되면서 본 모형의 성능을 확실하게 검증할 수 있었다.

본 연구 내용은 기존의 특허 문서로부터 토픽 분포 및 토픽-단어 분포의 시계열적 변동 정보로부터 다양한 기술의 경시적 변화를 도출할 수 있고, 이를 통해 특정 기술의 성숙도를 이해하거나 자사 및 경쟁사의 연구개발 추이를 파악하는 등 다양한 전략적 취지로 활용될 수 있으며, 미래 시점에 출현할 특허 문서를 예측하여 사람들로 하여금 기술적으로 진보된 태도를 취할 수 있도록 돕기 때문에 사회적으로 명확한 기여가 있다고 판단된다.

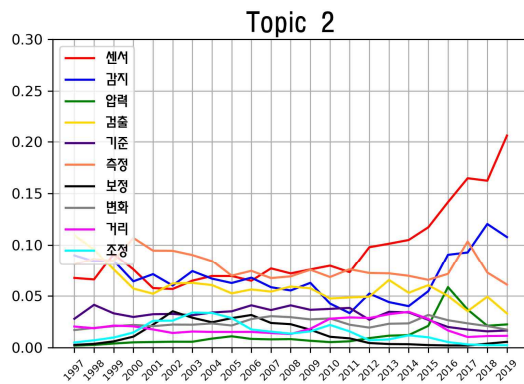
II. Figures and Tables



<그림 1> 주요 토픽 분포의 시계열적 변동



<그림 4> 미래 시점 토픽-단어 분포의 예측



<그림 2> 특정 토픽-단어 분포의 시계열적 변동

d_1 = [헤드 측면 구동부 모드 계조 소스 노출 발광 조사 발광 지지 ...]
 d_2 = [사용자 스테 모터 탄소수 반사 어부 드레인 두께 노출 플렉서블 유기 ...]
 d_3 = [패드부 면적 물질 버퍼 아이템 구조물 착용 픽셀 적색 어플리케이션 하부 ...]
 d_4 = [마이크 검사 임의 유닛 플레이트 렌즈 제조 추가 컴포넌트 표시장치 임의 ...]
 d_5 = [서브 시스템 휘도 전화 각도 헤테로시를 전압 연장 커버 통과 ...]

<그림 5> 예측된 토픽 및 토픽-단어 분포를 통해 생성된 미래 시점 문서

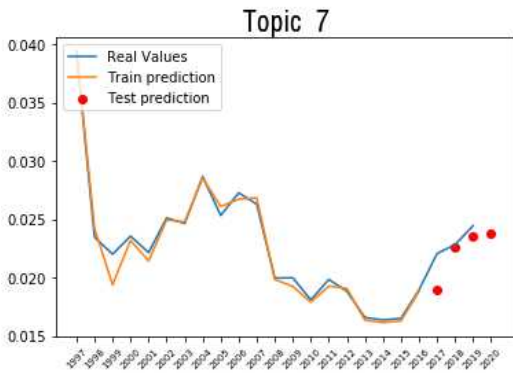
III. 참고문헌

Blei, David M., Andrew Y. Ng, and Michael I. Jordan. "Latent dirichlet allocation." Journal of machine Learning research 3.Jan (2003): 993-1022.

Blei, David M., and John D. Lafferty. "Dynamic topic models." Proceedings of the 23rd international conference on Machine learning, 2006.

Griliches, Zvi. "Patent statistics as economic indicators: a survey." R&D and productivity: the econometric evidence. University of Chicago Press, 1998. 287-343.

Kim, Gabjo, and Jinwoo Bae. "A novel approach to forecast promising technology through patent analysis." Technological Forecasting and Social Change 117 (2017): 228-237.



<그림 3> 미래 시점 토픽 분포의 예측

E3.5 BERT기반의 Context Sentence정보를 활용한 계약서 조항 분류 연구

전명준
연세대학교
공과대학

sangindong@yonsei.ac.kr

김우주
연세대학교
공과대학

wkim@yonsei.ac.kr

임영익
(주)인텔리콘
연구소

ceo@intellicon.co.kr

홍기현
(주)인텔리콘
연구소

kihyun.hong@intellicon.co.kr

정병택
(주)인텔리콘
연구소

jungbt@intellicon.co.kr

Abstract - 최근 인공지능의 기술 발전과 함께 법률 서비스에 적용하는 리걸테크의 산업이 확장되고 있다. 리걸테크의 분석 대상은 판례, 법령들과 같은 문서로 국한되었고 기업과 개인의 수요가 많은 계약서를 대상으로 한 연구는 제한적으로 진행되었다. 표준계약서로 검토할 경우 조항에 속해있는 세부 문장들에 대해서는 의미를 파악하기 어렵고, 표준계약서 없는 예외 조항들에 대해 유동적으로 대처하기 어려운 한계점이 존재한다. 법률 전문가들이 검토할 경우에도 많은 시간과 비용이 발생하는 문제점이 있다. 본 연구에서는 전문가의 지식이 집약된 자동화된 계약서 분석 방법론을 제안한다. 이를 위해, 계약서 세부 조항의 문장 분류를 진행하였으며 BERT기반의 문장 인코더를 사용해 길이가 긴 문서를 효과적으로 나타내는 새로운 문장 분류 방법론을 활용하였다. 본 연구의 결과로 계약서 검토에 대한 시간과 비용을 줄이고 법률 전문가의 지식을 활용할 수 있는 지능형 계약서 분석 시스템을 개발하였다.

Key Terms - Legal Tech, Contract Analysis, BERT, Sentence Classification

본 연구는 (주)인텔리콘 연구소에서 지원을 받아 수행한 연구임

I. 서론

최근 인공지능의 기술 발전과 함께 법률 서비스에 적용하고자 하는 시도가 활발하게 이루어지고 있다. 리걸테크(Legal Tech)는 법률(Legal)과 기술(Technology)의 합성어로 인공지능 기술을 이용해 법률서비스를 개선하고자 하는 산업을 의미하며, 미국과 유럽을 중심으로 리걸테크 산업이 확장되고 있다. 리걸테크를 구현할 수 있는 메커니즘으로는 규칙기반 시스템(Rule Based System)과 머신러닝 시스템이 있다(Thanaki, 2019). 규칙기반 시스템은 기존에 존재하는 지식 또는 규칙을 데이터셋에 적용하여 일정한 결과를 생성하거나 예측하는 시스템이다. 이 시스템은 인간의 전문적인 규칙이나 지식을 모방하고자 할 때 사용되며, 사전에 주어진 문법 규칙, 번역 규칙을 바탕으로 데이터셋을 학습해 나가는 방식이다. 머신러닝 시스템은 머신러닝 기술을 통해 확률 기반의 접근법을 활용하는 방식이다. 최근 딥러닝 기술이 발달하면서, 특히 단어의 의미를 효과적으로 표현하는 워드 임베딩(Word Embedding)방식인 CNN, LSTM과 같은 네트워크 모델을 효과적으로 활용할 수 있게 되었다. 최근에는 대용량 말뭉치를 사전 학습한 언어 모델인 BERT(Devlin et al., 2018)가 공개됨에 따라 현재

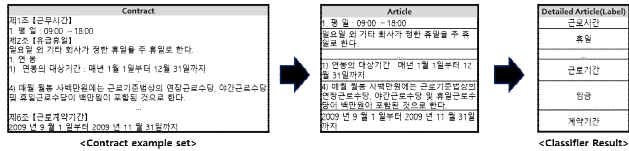
자연어처리에서 상당한 성능을 갖춘 인공지능 모델을 활용할 수 있게 되었다.

이러한 기술 발전과 함께 리걸테크의 분석 대상은 판례, 법령들과 같은 문서로 국한되었고 계약서를 대상으로 개체들의 정보 추출(Information Extraction)을 시도한 연구(Manor et al., 2019)가 진행되었으나 계약서에 있는 전체 조항을 대상으로 분석하는 연구는 진행되지 않았다. 따라서 이러한 한계점을 해결하기 위해 계약서 전체를 대상으로 분석할 수 있는 연구 개발의 필요성이 제기되고 있다.

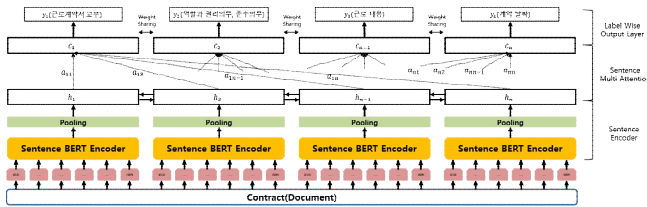
계약서는 기업 또는 개인이 상호 간에 합의한 계약 사항에 관하여 작성한 문서이다. 계약마다 목적과 규모는 다르지만 당사자 간의 이해관계가 얽혀있기 때문에 법률적인 검토가 필요하다. 이러한 계약서를 검토할 때는 필수적으로 검토해야 하는 사항들이 있기 때문에 표준계약서와 같은 양식으로 조항들을 파악하였지만 이러한 검토 방법은 다음과 같은 한계점을 가지고 있다. 첫 번째, 조항에 속해있는 세부 문장들에 대해서는 의미를 파악하기 어렵다. 두 번째, 계약당사자들에 따라서 조항에 명시된 구체화 정도가 다를 수 있기 때문에 중의적인 의미가 담긴 문장들을 검토하는데 한계가 있다. 세 번째, 표준계약서 이외의 법적인 분쟁 사항이 발생할 수 있는 예외 조항들을 발생할 수 있는 조항들에 대해 유동적으로 대처하기 어렵다. 설령, 위의 한계점을 전문가를 통해서 해결한다 하더라도 많은 시간과 비용이 발생하며 보통 계약 문서 한 건을 검토하는 데 1주일에서 최대 한 달 가량의 시간이 걸리는 문제점이 존재한다(Timmer et al., 2016). 따라서 본 연구에서는 이러한 문제점을 해결하기 위해 진화된 리걸테크의 기술을 활용하여 전문가의 지식이 집약된 새로운 계약서 분석 방법론을 제안한다.

본 연구에서 제안하는 계약서 분석 방법론은 <그림 1>과 같이 계약서를 문장 단위로 분리한 뒤, 계약서에서 필수적으로 검토되어야 할 사항들을 세부 조항(Label)으로 정의하여 검토 대상의 문장들을 분류한다. 분류 모델은 BERT를 활용하여 문장을 표현하고 검토 대상의 문장과 주변 문장과의 관계를 통해 문장을 효과적으로 나타내는 메커니즘을 사용하며, 제안하는 모델의 성능을 검증하기 위해 BERT를 포함한 3가지 모델을 베이스라인 모델로 선정하여 성능을 비교분석 하였다.

II. Figures and Tables



<그림 1> 제안 계약서 분석 방법론



<그림 2> 제안 문장 분류 방법론

<표 1> 근로계약서 문장 분류 성능 측정 결과

	Accuracy	Precision	Recall	F1-Score
CNN	73.87%	92.10%	71.78%	80.68%
LHAN	82.21%	91.38%	80.55%	85.62%
BERT	85.68%	87.91%	86.14%	87.02%
Proposed Model	86.64%	90.18%	86.25%	88.18%

III. 참고문헌

Chalkidis, I., Fergadiotis, M., Malakasiotis, P., Aletras, N., Androutsopoulos, I., “Extreme multi-label legal text classification: A case study in EU legislation” arXiv preprint arXiv:1905.10892 (2019).

Devlin, J., Chang, M. W., Lee, K., Toutanova, K., “BERT: Pre-training of deep bidirectional transformers for language understanding” arXiv preprint arXiv:1810.04805 (2018).

Kim, Y., “Convolutional neural networks for sentence classification” Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing, (2014), 1746-1751.

Liu, P., Qiu, X., Huang, X., “Recurrent neural networks for text classification with multi-task learning” arXiv preprint arXiv:1605.05101 (2016).

Manor, L., Li, J. J., “Plain English Summarization of Contracts” arXiv preprint arXiv:1906.00424 (2019).

Reimers, N., and Gurevych, I., “Sentence-bert: Sentence embeddings using Siamese bert-networks” arXiv preprint arXiv:1908.10084 (2019).

Sulea, O. M., Zampieri, M. Vela, M., Van Genabith, J., “Predicting the law area and decision of French supreme court cases” arXiv preprint arXiv:1708.01681 (2017).

Timmer, I., “Changing roles of legal: On the impact of innovations on the role of legal professionals and legal departments in contracting practice” Journal of Strategic Contracting and Negotiation, Vol.2, No.1-2(2016), 34-47.

Wan, L., Papageorgiou, G., Seddon, M., Bernardoni, M., “Long-length Legal Document Classification” arXiv preprint arXiv:1912.06905 (2019).

Wei, F., Qin, H., Ye, S., Zhao, H., “Empirical study of deep learning for text classification in legal document review” 2018 IEEE International Conference on Big Data, (2018), 3317-3320.

Yang, Z., Yang, D., Dyer, C., He, X., Smola, A., Hovy, E., “Hierarchical attention networks for document classification” Proceedings of the 2016 conference of the North American chapter of the association for computational linguistics: human language technologies, (2016), 1480-1489.



한국지능정보시스템학회

Korea Intelligent Information Systems Society