

설명가능 인공지능을 활용한 라이프로그 기반 치매 위험도 산정 방법에 관한 연구

천희웅 · 박혜연 · 이병주 · 홍수연 · 정준각[†]

한양대학교 산업융합학부

A Study on Dementia Risk Assessment Using Lifelog Data with Explainable Artificial Intelligence

Huiwoong Cheon · Hyeyeon Park · Byeongju Lee · Suyeon Hong · Junegak Joung

School of Interdisciplinary Industrial Studies, Hanyang University

This paper presents a novel method for accurately assessing dementia risk by utilizing explainable artificial intelligence (eXplainable AI) and lifelog data, which plays a crucial role in the early diagnosis of dementia. This study focuses on calculating dementia risk scores that not only derive the accuracy of the predictive model but also interpret the model's prediction results using SHAP values. This enables us to provide patients with clearer and more specific follow-up plans. The significant contribution of this study is that, based on the calculated scores, individuals with a high risk of dementia are promptly guided to undergo cognitive screening tests (CIST), allowing dementia treatment to commence at the optimal stage. By individually explaining the impact of each feature on the prediction results, SHAP values assist medical professionals in better understanding and utilizing the model's predictions.

Keywords: AI, CIST, SHAP Values, Predictive Model, Feature Impact Analysis

1. 서론

최근 고령화 사회로의 진입과 함께 치매 환자 수가 급격히 증가하고 있다. 수명 연장과 고령화는 세계적 현상이며, 선진 7개국(미국, 캐나다, 영국, 프랑스, 독일, 이탈리아, 일본)에서 더욱 현저하다(Central Dementia Center, n.d.). 우리나라의 경우, 전체 인구에서 65세 이상 노인 인구 비율은 2023년 18.4%이며, 2025년에는 20.3%로 초고령 사회가 될 전망이다(Statistics Korea, 2023).

고령화 사회 진입으로 인한 치매 발병률 증가로 인해 치매 환자의 조기 진단과 치료가 중요한 과제로 부상하고 있다. 치매는 환자 개인뿐만 아니라 그 가족과 사회 전체에 걸쳐 심각한 영향을 미치며, 이에 따른 의료비 부담 역시 급격히 증가하

고 있다. 2021년 기준으로 전국 65세 이상 추정 치매 환자 수는 약 892,002명에 달하며, 이들에 대한 국가 단위 치매 관리 비용은 최근 6년간 점차 증가하고 있으며, 치매 환자 수가 급증함에 따라 14.2조 원에서 약 18.7조 원으로 31.9% 증가했다(Central Dementia Center, 2022). 그러나 향후 고령 인구의 증가와 함께 치매 환자 수와 관리 비용도 급증할 것으로 예상된다. 실제로, 2040년에는 치매 관리 비용이 약 56.9조 원에 이를 것으로 추정되고 있다. 이는 현재 대비 약 3.3배 증가한 수치로, 국가 제정에 막대한 부담을 줄 수 있다(Central Dementia Center, 2021).

치매는 무엇보다 예방이 중요한 질환이다. 치매를 조기에 발견하고 지속적인 치료가 이루어진다면 치매 증상 악화를 지연시키는 효과가 있다고 알려져 있다(Kim, 2021). 치매의 발병을

이 논문은 한양대학교 교내연구지원사업과 한국연구재단의 지원을 받아 수행되었음 (HY-202300000003614, NRF-RS-2024-00344286).

[†] 연락저자 : 정준각 교수, 서울특별시 성동구 왕십리로 222, 제2공학관 503-2호 한양대학교 산업융합학부, Tel : 02-2220-2363, Fax : 02-2220-2363,

E-mail : june30@hanyang.ac.kr

2024년 8월 16일 접수; 2024년 9월 14일 게재 확정.

2년 지연시킬 경우 20년 후 치매 유병률이 80% 수준으로 낮아지고, 5년 지연시킬 경우 56% 수준으로 감소하는 것으로 확인되었다(Health Insurance Review & Assessment Service, 2020).

하지만 치매 조기 검진 수검률은 전반에도 미치지 못하는 것으로 나타났다. 2021년 치매 위험군으로 분류되는 65세 이상 인구수가 8,569,865명이었던 것에 비해(Statistics Korea, 2024), 같은 해 시행한 치매 조기 검진 선별검사의 수검자는 3,863,894명으로 45.09%의 비율만이 조기 검진을 수검하였다. 이는 보다 많은 인구가 조기 검진에 참여할 방안이 필요함을 시사한다. 데이터를 기반으로 인지장애 여부를 예측하고, 이를 활용하여 조기 검진 참여를 촉진한다면 수검률을 증가시킬 수 있을 것이다.

기존의 연구들은 주로 경도인지장애(Mild Cognitive Impairment; MCI) 예측을 목적으로 다양한 머신러닝 기법을 활용한 데이터 분석에 집중해 왔다(Hwang and Ha, 2023; Lee and Oh, 2021; Choi *et al.*, 2023; Lee *et al.*, 2018). 이러한 연구들은 기계학습 알고리즘을 적용하여 데이터를 분석하고 예측 모델을 개발하여, 주로 모델의 정확도를 계산하고 이를 바탕으로 성능을 평가하는 데에 중점을 두고 있다. 그러나 이러한 접근법에는 모델의 해석 가능성 및 실질적인 활용 방안에 대해 상대적으로 소홀히 다루고 있다는 한계점과, 모델의 설명력을 충분히 고려하지 않으려는 경향이 있다.

이에 본 연구에서는 기존의 예측 모델 개발 및 성능 평가 방법론을 고수하면서도, 모델의 해석 가능성을 높이는 새로운 접근법을 제시하고자 한다. 즉, 기존 연구의 한계를 극복하고자 설명가능 인공지능(explainable Artificial Intelligence; XAI)을 도입하여 라이프로그 데이터를 기반으로 한 ‘치매 위험도 점수 산정 방법’을 연구하였다. 본 연구는 예측 모델의 정확도를 도출함과 동시에 SHAP(SHapley Additive exPlanation) value를 활용하여 모델의 예측 결과를 해석하고(Lundberg and Lee, 2017), 이를 통해 환자들에게 보다 명확하고 구체적인 후속 방안을 제시할 수 있도록 ‘치매 위험도 점수 산정’에 초점을 두었다. 본 연구의 중요한 기여는 점수 산정을 토대로, 치매 위험도가 높은 사람들에게 ‘인지선별검사(Cognitive Impairment Screening Test; CIST)’를 신속히 받도록 유도함으로써, 최적의 단계에서 치매 치료를 시작할 수 있게 하는 것이다. 이를 통해 조기 진단과 예방에 기여할 수 있다. 선제적인 관리는 치매의 진행을 늦추거나 예방할 수 있는 중요한 요소이며, 본 연구는 이를 효과적으로 실현하고자 한다. SHAP value는 각 특징이 예측 결과에 미치는 영향을 개별적으로 설명해 줌으로써, 의료 전문가들이 모델의 예측 결과를 더욱 쉽게 이해하고 활용할 수 있도록 돕는다. 이러한 접근법은 단순히 예측 모델의 성능을 높이는 데에 그치지 않고, 예측 결과를 실질적인 의료 환경에서 효과적으로 활용할 수 있는 방안을 제시함으로써, 기존 연구들의 한계를 보완하고자 한다. 따라서 본 논문은 경도 인지장애 예측 모델의 개발 및 평가에 있어, 모델의 해석 가능성과 실질적인 활용 방안을 제시함으로써 기존 연구들의 제한

점을 극복하고자 한다. 이를 통해 경도인지장애의 조기 예측 및 효과적인 대응 방안을 마련하는 데에 기여하고자 한다.

본 연구는 AI-Hub(<https://www.aihub.or.kr/>)에서 제공하는 “치매 고위험군 웨어러블 라이프로그” 데이터로 예측 모델을 구축하고 치매 위험도 점수를 도출하는 과정을 다룬다. 2장에서는 기존 연구 대비 본 연구의 기여도를 구체적으로 설명하고자 한다. 3장에서는 설명가능 인공지능 기법을 적용해 모델의 해석과 SHAP value를 활용한 치매 위험도 점수 산정 방법론을 다룬다. 4장에서는 예측 모델 구축 결과 및 치매 위험도 점수를 도출한다. 이후 점수의 통계적 검증을 수행한다. 마지막으로, 5장에서는 본 연구의 의의와 개선 방안을 제시한다.

2. 기존 문헌

2.1 데이터 기반 치매 진단 연구

기존의 연구들은 주로 경도인지장애와 치매 진단을 위한 데이터 기반 접근법을 활용하여 다양한 머신러닝 기법으로 데이터 분석 및 예측 모델 개발에 중점을 두고 있다. 이러한 연구들은 주로 기계학습 알고리즘을 적용하여 대규모 데이터셋을 분석하고, 예측 모델의 정확도를 평가하며 성능을 검증하는 데 초점을 맞추고 있다.

전이학습과 히트맵 시각화를 활용한 선행 연구에서는 ResNet-101 모델을 사용하여 전이학습을 수행하고, Grad-CAM 기술을 활용하여 모델의 의사결정 근거를 시각화하는 기법을 제안했다(Hwang and Ha, 2023). 전이학습 모델의 정확도는 84.21%로 나타났으며, ResNet-101 아키텍처를 사용하여 훈련하고, SGDM Optimizer와 20회의 Epoch로 최적화했다. 또한 단면 MRI 데이터를 사용하여 치매 환자와 정상인을 구분하는 능력을 학습했다. OASIS-3 데이터셋을 활용하여 치매를 예측하는 모델을 제안하는 연구에서는 PCA(Principal Component Analysis)를 사용하여 데이터의 차원을 축소하고, 그래디언트 부스팅과 Stacking 등 다양한 머신러닝 모델을 적용하여 성능을 비교했다(Lee and Oh, 2021). 이 연구는 이전 연구들과는 달리 뇌 생체 데이터뿐만 아니라 참가자의 성별과 의료 정보 데이터 등을 활용하여 치매 예측에 더욱 관련성이 높은 특징을 찾아내는 모델을 제안했다. 다양한 머신러닝 모델의 성능은 데이터셋의 밸런스를 조정하지 않고 모델을 훈련시켰을 때 Stacking 0.77251, Extreme Gradient Boosting 0.77251, Deep Neural Network 0.63507, Voting 0.76777로 나타났다. 자동화된 기계학습(AutoML)과 라이프로그를 활용하여 인지기능 장애 예측 모형을 개발한 연구에서는 다양한 기계학습 알고리즘을 학습용 데이터에 탐색적으로 적용하고 검증하면서 가장 우수한 기계학습 알고리즘을 효과적으로 선정했다(Choi *et al.*, 2023). 해당 연구에서는 인지기능 장애 예측에 라이프로그 데이터의 활용 가능성을 보였으며, 사용된 모델 중 Voting Classifier, Gradient Boosting Classifier, Extreme Gradient Boosting, Light

Gradient Boosting Machine, Extra Trees Classifier, Random Forest Classifier 모형이 높은 예측 성능을 보였고, 수면 중 평균 호흡수와 평균 심박수가 가장 중요한 특성 변수로 확인되었다. 경도인지장애의 조기 진단을 용이하게 하기 위해 인간 행동의 잠재적 특징을 조사한 연구에서는 라이프로그 데이터를 기반으로 특정 특징을 추출한 후 인공 신경망을 사용하여 환자를 분류하였다. 그 결과, 라이프로그 기반 분류기는 건강한 대조군과 경도인지장애 환자를 구별할 수 있는 유의미한 결과(ROC-AUC 0.81)를 보였다(Lee *et al.*, 2018).

위와 같이 기존 연구에서는 XAI를 활용하여 치매 위험도 점수를 만드는 연구는 거의 없다. 대부분의 연구는 예측 모델의 성능 향상에만 집중하며, 모델이 왜 특정한 예측을 내리는지에 대한 설명을 제공하지 않는다. 이에 본 연구는 XAI 기술 중 SHAP value 분석을 통해 치매 위험도 점수를 산정하는 방법에 대해 제시하고자 한다.

2.2 라이프로그를 활용한 건강 진단 연구

라이프로그 데이터를 활용한 건강 진단 연구는 최근 들어 많은 주목을 받고 있다. 라이프로그 데이터는 개인의 일상생활에서 발생하는 다양한 활동과 생체 신호를 연속적으로 기록한 데이터로, 이를 분석함으로써 개인의 건강 상태를 모니터링하고 예측할 수 있는 가능성을 제공한다. 기존 연구들은 주로 이러한 라이프로그 데이터를 활용하여 운동량, 수면 패턴, 심박수 등의 생리적 신호를 분석하고, 이를 바탕으로 기본적인 건강 상태를 평가하는 데 집중해 왔다.

한국전자통신연구원(ETRI) 라이프로그 데이터를 활용하여 Stacking 앙상블 구조를 제안하는 연구에서는 수면의 질을 평가하였다. 사용된 데이터셋은 스마트폰의 IMU 및 GPS 데이터 등으로부터 얻은 생리 반응 신호와 사용자가 직접 입력한 행동, 환경, 감정 및 수면과 관련된 설문조사 결과이다(Lim *et al.*, 2023). 스마트폰의 센서를 사용하여 수집한 사용자 행동 패턴

데이터와 사용자가 입력한 식단 및 활동 정보 데이터를 기반으로 한 연구에서는 비만 예방 시스템을 제안하였다. 스마트폰의 가속도 센서로 사용자의 속도 변화를 감지하여 스마트폰의 움직임 정도를 파악한다. 이를 통해 사용자의 걸음 수와 땀 수를 측정하여 행동 패턴을 분석한다. 개인의 식습관, 행동 패턴, 운동량 등을 분석하고 이를 토대로 온톨로지를 활용하여 적합한 운동을 추천한다(Yoo *et al.*, 2016).

이러한 연구들은 라이프로그 데이터를 활용하여 개인의 건강 상태를 지속적으로 모니터링하고, 예방적인 건강관리를 가능하게 하는 중요한 기여를 하고 있다. 우리 연구는 라이프로그 데이터를 활용하여 보다 심각한 질병인 치매의 조기 진단을 목표로 한다. 본 연구에서는 라이프로그 데이터를 통해 수집된 다양한 생리적 및 행동적 데이터를 분석하고, SHAP value를 활용하여 각 데이터 특징이 치매 위험도에 미치는 영향을 설명한다. 이러한 접근법은 단순히 건강 상태를 모니터링하는 데 그치지 않고, 실질적인 의료 환경에서 효과적으로 활용할 수 있는 구체적인 방안을 제시함으로써, 기존 연구의 한계를 보완하고자 한다.

3. 연구 방법

3.1 데이터 수집 및 전처리

본 연구에서는 AI-Hub(<https://www.aihub.or.kr/>)에서 제공하는 “치매 고위험군 웨어러블 라이프로그” 데이터를 분석 대상으로 한다. 연구에 사용된 데이터는 반지 형태의 데일리 수면, 활동 데이터 수집기를 통해 착용자의 수면 데이터(수면 시간, 수면 효율, 수면의 깊이 등)와 활동 데이터(활동 시간, 운동 시간, 회복 시간 등)를 수집하였다. 데이터 수집 대상은 전문의의 병리진단을 바탕으로 55세 이상의 정상인지군(CN)(무증상치매(aAD) 포함), 전조증상(MCI), 치매(Dementia)인 174명을 대상으로 하였으며, 특히 알츠하이머병(AD) 고위험군을 우선으

Subjects Inclusion Criteria	Individuals who fully understand the purpose of the study and voluntarily provide written informed consent
	Individuals without severe hearing, vision, or language disabilities that affect the completion of neuropsychological tests or other assessments
	Individuals with no history or evidence of alcohol or substance use disorders, stroke, or central nervous system disease or injury
Subjects Exclusion Criteria	Individuals who cannot read or write
	Individuals with convulsive disorder (epilepsy), depressive disorder, bipolar disorder, schizophrenia, or substance use disorder
	Individuals with a history of brain-related conditions (stroke, brain infections, Parkinson's disease, multiple sclerosis, cerebrovascular disease, etc.)
	Individuals taking long-term psychiatric medications (for depression, alcohol dependence)
	Individuals diagnosed with disabilities that hinder their participation in research (blindness in one eye, hearing loss in one ear)
	Individuals who have experienced severe head trauma with loss of consciousness
	Cancer patients within three years post-treatment
	Individuals with a history of metal implants
	Individuals who do not consent to participate in the study

Figure 1. Criteria for Data Collection Subjects

로 포함하였다. 추가적인 대상자 획득 기준은 <Figure 1>과 같다. 또한, 개인정보를 제거하는 비식별화 처리가 이루어져 데이터의 독립성을 유지하였다. 수집된 데이터는 1인당 35~122 일 치가 기록되어 있어서 한 사람당 적어도 1달 이상의 생활 패턴이 포함된 데이터이다. 본 연구에서는 하루 단위로 기록된 라이프로그 데이터를 통해 치매 위험도 점수를 산정한다.

수집한 데이터에서 분석에 의미 없는 데이터 삭제와 모델링 가능한 형태로 변환 등의 전처리를 진행하였다. 라이프로그 측정 시작 시간과 측정 종료 시간 데이터 등 모든 데이터가 동일한 값을 가지는 데이터와 결측치 처리가 불가능한 5분 당 심박동 로그 데이터 등은 제거하였다. 수면 시작 시간과 수면 종료 시간 등 Time Stamp 포맷 데이터의 경우 24시에 대응하는 실숫값으로 변환하고, 수면 종료 시간에서 수면 시작 시간을 뺀 수면 시간 변수를 추가했다. 그 외에도 어노테이션이 기록된 활동 로그 시계열 데이터 등은 각 측정값의 개수에 해당하

는 변수를 추출하거나 표준편차, 분산, 평균, 사분위수 등 기술 통계량 변수를 추가하는 전처리를 진행했다. 구체적인 독립 변수 목록과 설명은 <Figure 2>와 같다.

종속 변수의 라벨링은 병리 진단 결과를 바탕으로 정상인지군(CN)과 인지장애군(MCI, Dem) 두 개의 클래스로 정의하였다. 이때 정상인지군을 0, 인지장애군을 1로 인코딩하였다. 정상인지군은 7,737건, 인지장애군은 4,446건으로 정상인지군:인지장애군 클래스 비율은 1.74:1로 나타났다.

3.2 치매 예측 모델 구축

(1) 모델 학습 및 검증 방법

모델의 학습과 검증은 K-fold 교차 검증(K-fold cross validation)을 활용하였다. K-fold 교차 검증은 모델의 일반화 성능을 평가하고 과적합을 방지하기 위해 필요하다. 단일 훈련/테

feature	description	feature	description
activity_average_met	Average daily MET	sleep_midpoint_time	Sleep midpoint time
activity_cal_active	Daily active calorie expenditure	sleep_onset_latency	Sleep onset latency
activity_cal_total	Total daily energy expenditure	sleep_period_id	Sleep period ID
activity_daily_movement	Daily movement distance	sleep_rem	REM sleep duration
activity_high	Duration of high-intensity activity	sleep_restless	Restlessness ratio during sleep
activity_inactive	Duration of inactivity	sleep_rmssd	Average heart rate variability during sleep
activity_inactivity_alerts	Number of inactivity alerts	sleep_score	Overall sleep score
activity_low	Duration of low-intensity activity	sleep_score_alignment	Sleep timing score
activity_medium	Duration of moderate-intensity activity	sleep_score_deep	Deep sleep score
activity_met_min_high	Daily MET for high-intensity activity	sleep_score_disturbances	Sleep disturbance score
activity_met_min_inactive	Daily MET for low-intensity activity	sleep_score_efficiency	Sleep efficiency score
activity_met_min_low	Daily MET for low-intensity activity	sleep_score_latency	Sleep latency score
activity_met_min_medium	Daily MET for moderate-intensity activity	sleep_score_rem	REM sleep score
activity_non_wear	Duration of non-wear time	sleep_score_total	Sleep contribution score
activity_rest	Duration of rest	sleep_temperature_delta	Skin temperature deviation during sleep
activity_score	Activity score	sleep_total	Total sleep duration
activity_score_meet_daily_targets	Daily target achievement score	activity_met_lmin_std	Standard deviation of 1-minute MET log per day
activity_score_move_every_hour	Hourly activity maintenance score	activity_met_lmin_variance	Variance of 1-minute MET log per day
activity_score_recovery_time	Recovery time score	activity_met_lmin_kurtosis	Kurtosis of 1-minute MET log per day
activity_score_stay_active	Activity maintenance score	activity_met_lmin_skewness	Skewness of 1-minute MET log per day
activity_score_training_frequency	Training frequency score	activity_met_lmin_mean	Mean of 1-minute MET log per day
activity_score_training_volume	Training volume score	activity_met_lmin_median	Median of 1-minute MET log per day
activity_steps	Average daily step count	activity_met_lmin_min	Minimum of 1-minute MET log per day
activity_total	Total activity duration (minutes)	activity_met_lmin_max	Maximum of 1-minute MET log per day
sleep_awake	Time awake during sleep period	activity_met_lmin_autocorrelation	Autocorrelation of 1-minute MET log per day
sleep_bedtime_end	Bedtime end (wake-up time)	activity_met_lmin_quantile_25	1-minute MET log Q1 per day
sleep_bedtime_start	Bedtime start	activity_met_lmin_quantile_50	1-minute MET log Q2 per day
sleep_time	Actual sleep duration	activity_met_lmin_quantile_75	1-minute MET log Q3 per day
sleep_breath_average	Average respiratory rate (breaths per minute)	activity_class_5min_count_1	Number of rest periods in 5-minute activity log per day
sleep_deep	Deep sleep duration	activity_class_5min_count_2	Number of inactive periods in 5-minute activity log per day
sleep_duration	Total time in bed	activity_class_5min_count_3	Number of low-intensity periods in 5-minute activity log
sleep_efficiency	Sleep efficiency	activity_class_5min_count_4	Number of medium-intensity periods in 5-minute activity log
sleep_hr_average	Average sleep heart rate per minute	sleep_hypnogram_5min_count_1	Number of deep sleep periods in sleep state log
sleep_hr_lowest	Lowest sleep heart rate per minute	sleep_hypnogram_5min_count_2	Number of light sleep periods in sleep state log
sleep_light	Light sleep duration	sleep_hypnogram_5min_count_3	Number of REM sleep periods in sleep state log
sleep_midpoint_at_delta	Sleep midpoint time delta	sleep_hypnogram_5min_count_4	Number of awakenings in sleep state log

Figure 2. Feature List

스트 분할 방식은 데이터 분할에 따라 성능이 달라질 수 있지만 K-fold 교차 검증은 전체 데이터 셋에 대한 모델의 성능을 더욱 정확하게 평가하고 과적합을 방지할 수 있다. 본 연구에서는 데이터 셋을 5개의 폴드로 나누어 교차 검증을 실시하였다.

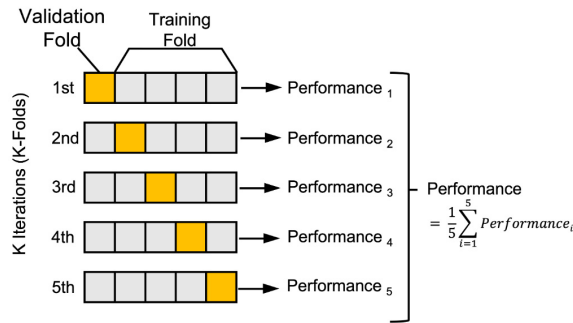


Figure 3. 5-fold Cross Validation

(2) 예측 모델 선정 및 중요 변수 선택

치매 위험도 점수 산정에 사용할 예측 모델의 선정에 위해 7가지(Logistic Regression, Decision tree, K-Nearest Neighbor, Support vector machine, Multi-Layer Perceptron, Random forest, LightGBM) 베이스 알고리즘으로 종속 변수 예측을 진행하였다. 베이스 알고리즘은 해석 가능한 머신러닝으로 온라인 리뷰 마이닝을 통해 서비스 특징의 중요도를 추정하는 프레임워크를 제안한 연구(Shin *et al.*, 2024)를 참조하여 높은 성능을 보이는 4가지(Decision tree, Multi-Layer Perceptron, Random forest, LightGBM) 모델을 선정했고, 분류(Classification) 문제에서 기준 모형(baseline)으로 많이 사용되는 Logistic Regression, K-Nearest Neighbor, Support vector machine을 추가로 선정했다. Logistic Regression은 선형 회귀 모형에서 발전된 이진 분류 모델로, 독립 변수와 종속 변수 간의 관계를 설명하기 위해 로짓 변환을 사용해 분류를 진행하는 모델이다(Hosmer and Lemeshow, 2000). Decision Tree(이하 DT)는 데이터의 속성에 따라 분할 규칙을 반복적으로 적용하여 if-else 구조의 트리를 형성해 데이터를 분류하는 모델이며(Quinlan, 1986), K-Nearest Neighbor는 예측 시 가장 가까운 k개의 이웃 데이터 포인트의 레이블을 다수결로 결합해 예측을 수행하는 비모수적 모델이다(Cover and Hart, 1967). Support Vector Machine은 데이터를 고차원 공간으로 변환한 후, 최적의 초평면을 찾아 데이터를 분류하는 모델이다(Cortes and Vapnik, 1995). Multi-Layer Perceptron은 다층 신경망 구조로 이루어진 모델로, 은닉층을 통해 비선형 문제를 학습하고 예측하는 데 사용되며, 역전파 알고리즘을 활용해 가중치를 학습하는 모델이다(Rumelhart *et al.*, 1986). Random Forest는 여러 개의 DT를 앙상블하여 예측을 수행하는 모델로, 개별 DT의 단점을 보완하고 일반화 성능을 높이는 데 유리한 모델이다(Breiman, 2001). LightGBM은 그라디언트 부스팅(Gradient Boosting) 알고리즘을 기반으로 한 모델로, 히스토그램 기반 학습 방식을 도입해 대규모 데이

터에서의 학습 속도와 성능을 크게 향상시킨 모델이다(Ke *et al.*, 2017).

각 모델의 학습을 위해 먼저 주어진 데이터 셋을 5개의 폴드로 나누고, 각 폴드마다 훈련 셋과 검증 셋으로 분할했다. 그 다음, 훈련 셋을 사용하여 각 분류 모델을 학습시켰다. 학습된 모델을 사용하여 검증 셋에 대한 예측을 수행했고, 예측 결과를 기반으로 성능을 평가하였다.

이후 ROC-AUC 점수를 기준으로 가장 높은 성능을 보여준 모델을 선정하였다. 본 연구의 예측 모델은 이진 분류 모형이므로, 모델의 전반적인 분류 능력을 다양한 임계값에서 평가할 수 있는 ROC-AUC 점수를 기준 지표로 삼았다. 해당 지표는 ROC(Receiver Operating Characteristic) 곡선의 하단 면적에 해당하는 AUC(Area Under the Curve)를 뜻하며, 본 연구에서는 연구 목적에 부합하도록 인지장애군 클래스에 해당하는 ROC 곡선을 사용하였다. ROC-AUC는 특히 클래스 불균형이 있는 데이터 셋의 이진 분류 문제에서 안정적인 성능 평가가 가능하다.

선정된 모델을 기반으로 Feature selection 기법 중 전진 선택법(Forward Selection)을 활용하여 중요 변수를 선택하였다. SHAP Importance를 이용하여 전체 특성별 중요도를 평가하였고, 중요도가 높은 변수부터 1개씩 추가하면서 예측 모델을 구축하였다. 구축된 예측 모델 중 가장 높은 성능을 보일 때의 독립 변수들을 중요 변수로 선택하여 기존보다 예측 성능을 향상시켰다.

(3) 하이퍼 파라미터 튜닝

마지막으로, 선정된 모델과 변수를 토대로 하이퍼 파라미터 튜닝을 통해 예측 모델의 성능을 최적화하여 최종 모델을 구축하였다. 하이퍼 파라미터는 학습 과정에서 모델이 어떻게 학습할지 결정하는 상수로, 모델 자체의 학습을 통해 결정되는 파라미터(가중치 등)와는 구별된다. 하이퍼 파라미터의 예시로, 결정 트리 기반의 그라디언트 부스팅(Gradient Boosting) 알고리즘을 구현한 LightGBM에는 하나의 트리에서 가질 수 있는 최대 리프 노드 개수를 뜻하는 'num_leaves'와 리프 노드에 들어가는 최소 데이터 수를 뜻하는 'min_data_in_leaf' 등이 있다.

본 연구에서는 그리드 서치(Grid Search)를 사용하여 하이퍼 파라미터 튜닝을 진행하였다. 그리드 서치는 미리 정의된 값의 조합을 모두 시도하여 최적의 조합을 찾는 튜닝 방법이다. 먼저 임의의 범위를 선정하여 그리드 서치를 수행한 후, ROC-AUC 점수를 기준으로 교차 검증을 통해 우수한 하이퍼 파라미터값을 구한 다음 근처의 범위를 새로운 후보값으로 선정하여 다시 그리드 서치를 수행하는 과정을 반복하였다. 4.1의 최종 하이퍼 파라미터 중 'num_leaves'의 경우, 먼저 100-1000 범위를 100 간격으로 탐색, 300이라는 값을 얻은 다음 220-380 범위를 20 간격으로 탐색, 320이라는 값을 얻은 다음 마지막으로 310-330 범위를 10 간격으로 탐색하여 최종값을 결정하였다.

3.3 예측 모델 해석 기반 치매 위험도 도출

최종 선정된 블랙박스 모델의 예측 결과를 해석하기 위해 SHAP(SHapley Additive exPlanation) value를 활용하여 Feature 별 기여도를 분석하고, 이를 바탕으로 치매 위험도 점수를 도출하는 새로운 방법론을 제시하였다. 본 방법론은 두 단계로 구성된다.

(1) SHAP value 도출

SHAP value는 게임 이론의 Shapley 값을 기반으로 하여 모델의 예측 결과에 각 Feature가 기여한 정도를 정량화하는 방법이다. SHAP value는 각 Feature가 예측 결과에 미치는 영향력을 개별적으로 계산한 후, 이를 모든 Feature의 조합에 대해 평균화하여 도출된다. 이는 다음의 수식으로 표현된다.

Equation 1. SHAP value

$$\phi(i) = \sum_{S \subseteq N/i} \frac{|S|!(|M|-|S|-1)!}{|M|!} (v(S \cup i) - v(S))$$

$\phi(i)$: 특정 Feature i 에 대한 SHAP value

N : 전체 Feature 개수

S : Feature의 부분 집합, 예측 모델에서 고려하는 Feature의 조합

$v(S)$: S 에 의한 예측 결과, 특정 S 가 주어졌을 때 모델의 예측값

$v(S \cup i)$: S 와 Feature i 를 포함한 조합에 의한 모델의 예측값

$|S|$: S 의 크기, S 에 포함된 Feature 개수

$|S|!(|M|-|S|-1)!$: 모든 가능한 조합에 대해 평균을 내는데 사용되는 가중치

전체 SHAP value로 도출되는 SHAP matrix는 모델이 예측을 수행한 후에 각 Feature가 해당 예측에 기여한 정도를 설명하는 값으로 구성된 데이터이다. 예측 모델을 구성하기 위해 모델이 학습하는 입력인 독립 변수와 예측해야 하는 목표인 종속 변수로 이루어진 원시 데이터 셋과는 다르게, SHAP matrix는 모델의 예측 결과에 대한 설명을 제공하는 값들로 이루어

져 있으므로, 이를 이용하여 예측에 대한 사후 해석 가능성을 높일 수 있다.

(2) 치매 위험도 산정

치매 위험도 점수를 도출하기 위해, 하루 단위로 계산된 SHAP value 중 양의 값을 합산하여 개별 피험자의 치매 위험도를 산출한다(Equation 2). 양의 SHAP value는 해당 Feature가 치매 위험도를 증가시키는 방향으로 기여한 정도를 의미하며, 이를 통해 피험자의 치매 위험도를 평가할 수 있다. 이러한 접근 방식은 개별 Feature가 치매 위험도에 미치는 양의 영향을 반영한 위험도 점수 산출을 가능하게 한다.

정상인지군과 인지장애군으로 분류된 피험자들에 대해, 각 그룹의 치매 위험도 점수 평균을 계산하여 그룹 간의 차이를 분석하였다. 이 분석을 통해 제안된 방법론이 치매 발병 소지를 정량적으로 평가하는 데 유의미한 정보를 제공함을 확인할 수 있다.

이를 통해 본 방법론에서는 SHAP value를 활용하여 Feature의 기여도를 정량화하고, 이를 통해 치매 위험도 점수를 도출하는 새로운 접근 방식을 제시한다.

Equation 2. Dementia Risk Score

$$DRS_i = \sum_{j=1}^n s_{ij}, \quad s_{ij} = \begin{cases} 0, & \text{if } x_{shap,ij} < 0 \\ x_{shap,ij}, & \text{if } x_{shap,ij} \geq 0 \end{cases}$$

DRS_i : 일일 데이터 i 에 대한 치매 위험도 점수

s_{ij} : 일일 데이터 i 의 j 번째 라이프로그 특성에 대한 라이프로그 위험 점수

n : 라이프로그 특성 개수

$x_{shap,ij}$: 일일 데이터 내 j 번째 라이프로그 특성에 대한 SHAP value

4. 연구 결과

4.1 치매 예측 모델 구축

예측 모델의 선정을 위해 7가지 알고리즘으로 종속 변수를 예측하였다. 각 알고리즘별 성능 비교 결과는 <Table 1>과 같

Table 1. Comparison of Performance Between Prediction Models

Model	Accuracy	ROC-AUC	Precision	Recall	F1-macro
LightGBM	0.8262	0.9010	0.8276	0.7904	0.8025
Random forest	0.8055	0.8835	0.8325	0.7491	0.7659
Decision tree	0.7041	0.6806	0.6808	0.6806	0.6807
K-Nearest Neighbor	0.6572	0.6595	0.6229	0.6111	0.6136
Mlti-Layer Perceptron	0.5953	0.6348	0.6188	0.5634	0.5142
Support vector machine	0.6393	0.6249	0.6609	0.5083	0.4114
Logistic regression	0.6457	0.6067	0.6113	0.5331	0.4830

다. 모델별 예측 성능 비교 결과 LightGBM 모델이 Accuracy 82.62%, ROC-AUC 0.9010, F1-macro 80.25%로 가장 좋은 성능을 보였다. 해당 모델에 대해 전진 선택법을 통한 Feature selection을 진행하였다. SHAP Importance를 기준으로 상위 40개의 변수를 사용할 때 ROC-AUC 0.9037로 가장 높은 예측 성능을 보였다. Feature selection 진행 과정에 대한 성능 평가 그래프는 <Figure 4>와 같다. 마지막으로, <Table 2>와 같이 하이퍼파라미터 튜닝을 진행하였을 때 ROC-AUC 0.9492로 가장 높은 성능을 보였다. 최종 예측 모델의 ROC 곡선은 <Figure 5>와 같다.

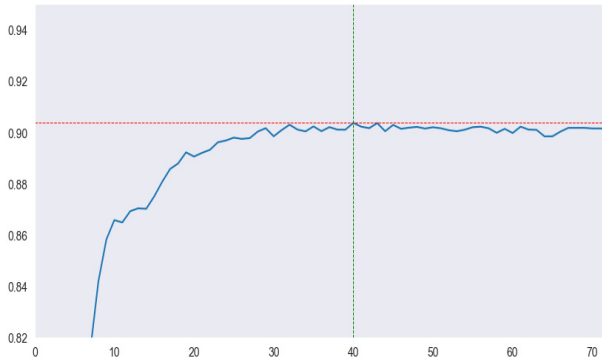


Figure 4. Performance Improvement with Forward Feature Selection

Table 2. LightGBM Hyper Parameter

Parameter	Value
min_data_in_leaf	41
num_leaves	330
n_estimators	1000
learning_rate	0.08

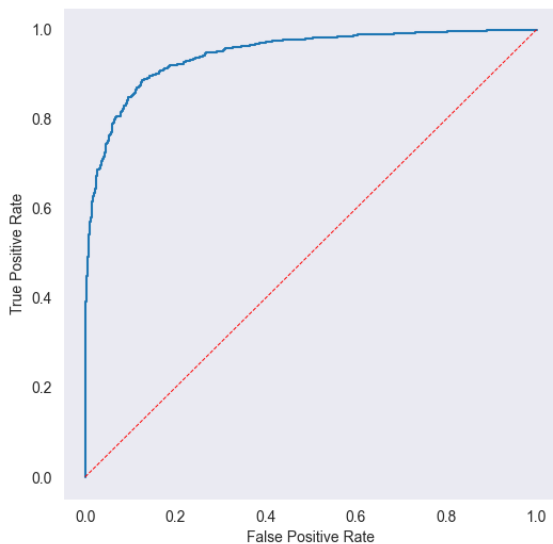


Figure 5. ROC curve of LightGBM

4.2 예측 모델 해석 및 치매 위험도 점수 도출

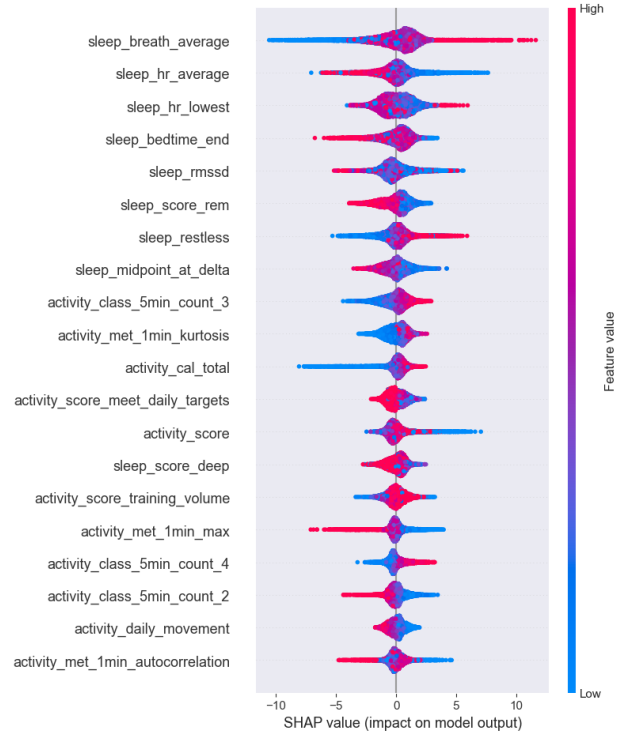


Figure 6. SHAP Summary Plot

예측 모델의 해석과 치매 위험도 점수 산정을 위해 SHAP value를 도출하였다. <Figure 6>는 각 변수에 대해 SHAP value가 어떠한 분포로 나타나는지 시각화한 그림이다. X축은 각각의 데이터에서 변수들의 SHAP value이다. 0을 기준으로 분포 범위가 왼쪽일수록 중속 변수에 대한 음의 영향력이 크고 오른쪽일수록 양의 영향력이 크다고 해석할 수 있다. 이때 색상은 변수의 값으로 빨간색에 가까울수록 큰 값, 파란색에 가까울수록 작은 값을 의미한다. Y축은 feature importance 순으로 나열한 것이다. 즉, 상위에 명시된 변수일수록 중요도가 높다고 해석할 수 있다.

분석 결과 전반적으로 수면 관련 변수가 주요한 영향을 미치는 것으로 확인되었다. 그중 중요도가 가장 높은 수면 시 분당 평균 호흡 수(sleep_breath_average)는 작은 값일 때 SHAP value가 낮고 큰 값일 때 SHAP value가 높다. 이는 수면 시 분당 평균 호흡 수가 적으면 정상인지군에 가깝고, 많으면 인지장애군에 가깝다고 해석할 수 있다. 행동 관련 변수 중에는 하루간 5분 당 활동 로그 중 낮은 강도 활동 개수(activity_class_5min_count_3)가 가장 높은 중요도를 보였다. 해당 변수도 작은 값일 때 SHAP value가 낮고 큰 값일 때 SHAP value가 높지만, 상대적으로 영향력은 모자란 것으로 나타났다. 이는 하루간 5분 당 활동 로그 중 낮은 강도 활동 개수가 적으면 정상인지군에 가깝고, 많으면 인지장애군에 가깝지만, 큰 영향을 주지는 않는다고 해석할 수 있다.

본 연구에서 제시한 치매 위험도 점수는 개별 변수가 치매

위험도에 미치는 양의 영향을 반영한다. 각 군 별 치매 위험도 점수를 산출한 결과 정상인지군의 경우 최소 1.06, 최대 24.99, 평균 7.59로 나타났으며 인지장애군의 경우 최소 1.79, 최대 31.28, 평균 15.71로 나타났다. 중요도가 높은 3개 변수를 보면 수면 시 분당 평균 호흡 수(sleep_breath_average)가 많을 때, 수면 시 분당 평균 심박동 수(sleep_hr_average)가 적을 때, 수면 시 분당 낮은 심박동 수(sleep_hr_lowest)가 많을 때 치매 위험도 점수가 높게 나오는 것을 확인할 수 있다.

생활 요인에 따른 치매 위험을 연구한 기존 연구들의 경우 대부분 포괄적인 범주의 요인 혹은 의학적 판단을 기준으로 하고 있고, 구체적인 생활 요인과 치매 발병 위험 간의 명확한 연관성에 대한 연구는 부족한 상태이다. 불면증 및 비특이적 수면 문제를 포함한 전반적인 수면 장애가 모든 원인 치매의 발병 위험을 높이는 것을 확인한 연구(Shi et al., 2018)와 수면 장애로 인한 전신 염증 증가가 β -아밀로이드 부담을 증가시켜 알츠하이머의 위험을 높일 수 있다는 연구(Irwin and Vitiello, 2019), 일일 총 신체 활동량이 많을수록 알츠하이머 발병 위험이 낮아지는 것을 확인한 연구(Buchman, 2012) 등 대부분이 넓은 범주의 생활 요인에 대해서만 다루고 있을 뿐이다. 그에 반해 본 연구에서는 SHAP 분석을 통해 치매 위험에 영향을 주는 더욱 구체적인 생활 요인들을 확인할 수 있고, 그중 수면 시 분당 평균 호흡 수(sleep_breath_average)는 가장 많은 영향을 미치며, 해당 요인이 큰 값일 때 치매 위험도가 높다는 등의 추가 해석도 확인할 수 있다.

4.3 검증

치매 위험도 점수를 도출하기 위해 정상인지군과 인지장애군을 구분하는 예측 모델을 사용하고, 독립 변수들의 영향력을 기반으로 정량적인 치매 위험도 점수를 산출했다. 각 군 별 치매 위험도 점수의 히스토그램은 <Figure 7>과 같다.

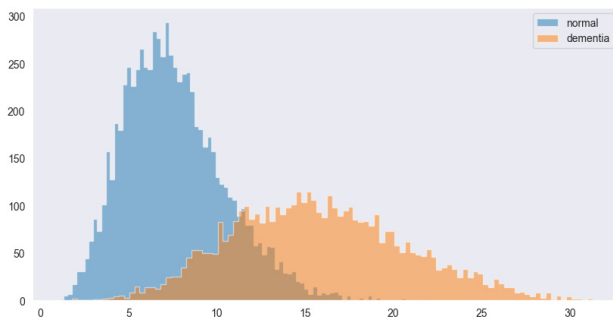


Figure 7. Histogram of Dementia Risk Score

하지만, 이러한 방법론이 치매 여부를 정확히 예측할 수 있는지 엄밀하게 증명하려면 실제 치매 위험도 점수가 있어야 하지만, 실제로 그러한 점수는 존재하지 않는다.

따라서, 이를 간접적으로 검증하는 방법으로 실제 레이블

(정상인지군, 인지장애군)을 기반으로 산출된 치매 위험도 점수를 두 그룹에 대해 구한 후, 정상인지군의 평균 치매 위험도 점수를 기준으로 설정한 다음 인지장애군의 치매 위험도 점수가 이 기준값보다 유의미하게 큰지 평가했다. 인지장애군의 평균이 기준값보다 유의미하게 클 경우, 제안된 치매 위험도 점수가 치매 여부를 예측하는 데 유효하다고 판단할 수 있다.

(1) 가설 설정:

- 귀무가설(H_0): 인지장애군(u_2)의 치매 위험도 점수의 평균은 정상인지군(u_1)의 치매 위험도 점수의 평균보다 작거나 같다($u_2 \leq u_1$).
- 대립가설(H_1): 인지장애군(u_2)의 치매 위험도 점수의 평균은 정상인지군(u_1)의 치매 위험도 점수의 평균보다 크다($u_2 > u_1$).

(2) t-검정 수행:

- 단일 표본 t-검정(one-sample t-test)은 특정 집단의 평균이 주어진 기준값과 유의미하게 다른지 평가하기 위해 주로 사용된다. 이를 통해 인지장애군의 치매 위험도 점수가 정상인지군의 평균보다 큰지를 확인할 수 있다.
- 단측 검정을 사용하여 p-value를 계산했다. p-value는 관측된 t-통계량이 귀무가설 하에서 우연히 발생할 확률을 나타내며, 일반적으로 p-value가 0.05보다 작을 경우 관측된 데이터가 귀무가설과 일치하지 않고 대립가설을 지지한다고 간주한다.

(3) 결과

단일 표본 t-검정 결과, p-value가 0에 수렴하여 귀무가설을 기각한다($p < 0.05$). 즉, 인지장애군의 치매 위험도 점수가 정상인지군의 평균보다 유의미하게 크다는 결론을 도출할 수 있다. 이를 통해 본 연구에서 제안한 치매 위험도 점수 산정 방법이 두 집단 간의 인지 상태 차이를 반영하며, 치매 여부를 예측하는 데 유의미한 지표임을 입증할 수 있다.

5. 결 론

본 연구에서는 설명 가능한 인공지능(XAI)을 활용한 라이프로그 학습 데이터 모델에 대해 소개하고, 이를 기반으로 한 치매 위험도 점수를 산정하였다. 사용한 데이터는 “치매 고위험군 웨어러블 라이프로그 데이터”로 55세 이상의 정상, 전조증상 및 치매를 진단받은 174명을 대상으로 하였다.

선행연구와 달리 본 연구는 라이프로그를 활용하여 치매 예측을 했고 SHAP value를 통한 치매 위험도 점수를 정량적으로 산정했다. 이를 통해 웨어러블 기기를 착용하는 것만으로도 누구나 치매 위험도를 파악할 수 있게 했다.

이러한 연구 결과를 바탕으로 위험도가 높은 사람들에게

“인지선별검사(CIST)” 신속히 받도록 유도하고 조기 진단과 예방에 기여할 수 있다.

그러나 본 연구는 데이터를 수집한 대상의 수가 174명으로 적고, 정상인지군과 인지장애군의 비율이 1.74 : 1로 불균형하다는 한계점을 가지고 있다. 이 때문에 사람 단위가 아닌 하루 단위의 라이프로그 데이터로 치매 예측 모델을 만들었다. 따라서, 하루 단위로 산정된 치매 위험도 점수로 특정 사람이 치매임을 진단하기 위해서는 추가적인 연구를 통한 기준 수립이 필요하다.

마지막으로, 다른 연구에서 시도하지 않았던 점수 산정 방법의 구현 가능성은 아주 높다고 사료된다. 이러한 점수 산정 방법을 통해 치매 고위험군을 조기에 식별하고, 예방 및 치료에 필요한 정보를 제공할 수 있을 것이다. 향후 누적된 데이터가 많아지면 더욱 신뢰도 높은 치매 위험도 점수를 산정할 수 있을 것으로 예상된다.

참고문헌

- Buchman, A. S., Boyle, P. A., Yu, L., Shah, R. C., Wilson, R. S., and Bennett, D. A. (2012), Total daily physical activity and the risk of AD and cognitive decline in older adults, *Neurology*, **78**(17), 1323-1329.
- Breiman, L. (2001), Random forests, *Machine Learning*, **45**(1), 5-32.
- Central Dementia Center (n.d.), Dementia Dictionary. https://www.nid.or.kr/info/diction_list2.aspx?gubun=0201.
- Central Dementia Center (2022), Status of Dementia in Korea 2021(p. 31). https://www.nid.or.kr/info/dataroom_view.aspx?bid=243.
- Central Dementia Center (2023), Status of Dementia in Korea 2022(pp. 14, 31). https://www.nid.or.kr/info/dataroom_view.aspx?bid=257.
- Choi, H. C., Yoon, C. H., and Lee, S. B. (2023), Cognitive Impairment Prediction Model Using AutoML and Lifelog, *Journal of the Korea Society of Computer and Information*, **28**(11), 53-63.
- Cortes, C. and Vapnik, V. (1995), Support-vector networks, *Machine Learning*, **20**, 273-297.
- Cover, T. and Hart, P. (1967), Nearest neighbor pattern classification, *IEEE Transactions on Information Theory*, **13**(1), 21-27.
- Health Insurance Review & Assessment Service (2020), One in Ten Elderly Individuals Has Dementia, Early Screening for Prevention is Essential (p. 9). <https://www.hira.or.kr/bbsDummy.do?brdBltno=10146&brdScnBltno=4&pageIndex=1&pgmid=HIRAA020041000100#none>.
- Hosmer, D. W. and Lemeshow, S. (2000), *Applied Logistic Regression* (2nd ed.). John Wiley & Sons, Inc.
- Hwang, H. and Ha, M. S. (2023), Explainable Dementia Prediction Diagnosis Using Transfer Learning and Heatmap Visualization, *Proceedings of the Korean Institute of Electrical Engineers Conference*, Jeju, 2520-2521.
- Irwin, M. R. and Vitiello, M. V. (2019), Implications of sleep disturbance and inflammation for Alzheimer's disease dementia, *The Lancet Neurology*, **18**(3), 296-306.
- Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., ... and Liu, T. Y. (2017), LightGBM: A highly efficient gradient boosting decision tree, *Advances in Neural Information Processing Systems*, **30**, 3146-3154.
- Kim, K. H. (2021), The Shock of Alzheimer's Drug Failure: Early Diagnosis of Mild Cognitive Impairment is Currently the Best Way to Delay Dementia, Chosunilbo, https://www.chosun.com/special/future100/fu_general/2021/01/18/PYULBOEVMZDNZDEA33ZXUSUPKU/.
- Lee, S., Kang, W., and Moon, C. (2018), Lifelog-Based Classification of Mild Cognitive Impairment Using Artificial Neural Networks, *2018 International Conference on Electronics, Information, and Communication (ICEIC)*, Honolulu, 1-2.
- Lee, T. I. and Oh, H. Y. (2021), Dementia Prediction Model based on Gradient Boosting, *Journal of the Korea Institute of Information and Communication Engineering*, **25**(12), 1729-1738.
- Lundberg, S. M. and Lee, S. I. (2017), A Unified Approach to Interpreting Model Predictions, *Advances in Neural Information Processing Systems*, 30.
- Lim, Y. H., Lee, M. H., Park, Y. I., Jeong, Y. I., Kang, E. S., Lee, C. W., ... and Han, H. W. (2023), Multi-layer Stacked Ensemble Models for Prediction of Sleep Quality Using Lifelog Data, *Proceedings of the Korean Institute of Information Scientists and Engineers Conference*, Jeju, 2037-2039.
- Quinlan, J. R. (1986), Induction of decision trees, *Machine Learning*, **1**(1), 81-106.
- Rumelhart, D. E., Hinton, G. E., and Williams, R. J. (1986), Learning representations by back-propagating errors, *Nature*, **323**(6088), 533-536.
- Shi, L., Chen, S. J., Ma, M. Y., Bao, Y. P., Han, Y., Wang, Y. M., ... and Lu, L. (2018), Sleep disturbances increase the risk of dementia: A systematic review and meta-analysis, *Sleep Medicine Reviews*, **40**, 4-16.
- Shin, J., Joung, J., and Lim, C. (2024), Determining directions of service quality management using online review mining with interpretable machine learning, *International Journal of Hospitality Management*, **118**, 103684.
- Statistics Korea (2023), 2023 Statistics on the Aged Population(p. 8). https://kostat.go.kr/board.es?mid=a10301010000&bid=10820&act=view&list_no=427252.
- Statistics Korea (2023), Key Demographic Indicators (Sex Ratio, Population Growth Rate, Population Structure, Dependency Ratio, etc.). https://kosis.kr/statHtml/statHtml.do?orgId=101&tblId=DT_1BPA002&conn_path=I2.
- Yoo, M. H., Cho, J. H., Byun, S. H., and Lee, S. W. (2016), System for Preventing Obesity Based on User Lifelog, *Proceedings of Symposium of the Korean Institute of communications and Information Sciences*, Jeju, 606-607.

저자소개

천희용 : 한양대학교 산업융합학부 정보공학전공 학사 과정에 재학 중이다. 관심 연구 분야는 설명가능 인공지능을 활용한 사람 중심 설명 지원, 설명가능 인공지능 기법의 사용자 맞춤화이다.

박혜연 : 한양대학교 산업융합학부 정보공학전공 학사 과정에 재학 중이다. 관심 연구 분야는 설명가능 인공지능을 활용한 예측 시스템 개발, 품질 데이터 분석이다.

이병주 : 한양대학교 산업융합학부 정보공학전공 학사 과정에 재학 중이다. 관심 연구 분야는 데이터 기반의 프로세스 혁신을 통한 전통 산업군의 업무 효율성 개선이다.

홍수연 : 한양대학교 산업융합학부 정보공학전공 학사 과정에 재학 중이다. 관심 연구 분야는 Data-Driven UX 기반 고객 경험 데이터 분석이다.

정준각 : 포항공과대학교 산업경영공학과에서 2013년 학사, 2019년 석박사 통합 학위를 취득하였다. 일리노이 대학교 어바나-샴페인과 울산과학기술원 산업공학과에서 박사후연구원으로 근무했으며, 현재 한양대학교 산업융합학부에서 조교수로 재직 중이다. 관심 연구분야는 텍스트 마이닝, 품질 데이터 분석, 설명가능 인공지능 응용이다.