

논문 2022-2-15 <http://dx.doi.org/10.29056/jsav.2022.12.15>

랜덤 포레스트를 이용한 스마트폰 사용자의 행위 기반 분류의 글로벌 해석을 위한 시각화

김영인*†

Visualization for Global Interpretation of Activity Based Classification of Smartphone Users using Random Forest

Young-In Kim*†

요 약

모바일 어플리케이션에서 편리하고 지속적이면서 에너지 효율이 높은 사용자 분류 기술의 중요성이 증가하고 있으며, 다양한 사용자 움직임을 기반으로 성능이 좋은 사용자 분류 방법에 대한 연구가 활발히 이루어지고 있다. 하지만 현재 연구되고 있는 사용자 분류 기법은 사용한 기계 학습 모델이 크고 복잡하여 분류 결과를 해석하여 설명하기는 어려운 실정이므로, 이를 개선하여 연구자가 사용자 분류 결과에서 유사한 사용자와 아닌 사용자의 원인을 설명할 수 있는 방안에 대한 연구가 필요하다. 이를 위하여 본 논문에서는 하이퍼파라미터를 최적화하여 개발한 랜덤 포레스트 분류 모델을 적용하여 사용자를 분류하고, 그 결과에서 분류 결정을 하게 된 글로벌 해석을 시각화하여 설명으로 제공하고자 설명가능 AI(XAI)의 SHAP(Shapley additive explanation)을 사용하여 제시한다. 실험을 통하여 실험 데이터에 있는 특징 중, 어떤 특징들이 사용자 분류에 가장 크게 기여하는 것인지 요약한 결과 및 정분류된 사용자와 오분류된 사용자의 인스턴스를 해석한 결과를 제시하였다. 그 결과, 랜덤 포레스트 분류 모델이 사용자 분류 과정에서 사용자를 분류한 이유를 해석할 수 있음을 보였으며, 사용자 행위 센서 데이터를 사용한 사용자 분류의 설명 가능한 기법으로 SHAP을 사용할 수 있음을 제시하였다.

Abstract

The importance of convenient, continuous and energy-efficient user classification technology is increasing in mobile applications, and research on user classification methods with good performance based on various user movements is being actively conducted. However, the user classification techniques currently being studied have a large and complex machine learning model, so it is difficult to interpret and explain the classification result. So research is needed to study how to improve this so that researchers can explain the causes of similar users and dissimilar users in the user classification results. In this paper, we classify users by applying the Random Forest classification model developed by optimizing the hyperparameters, and use SHAP(Shapley additive explanation) of explainable AI(XAI) to visualize and provide the global interpretation for the classification decision from the results. Through the experiment, the results of summarizing which of the features in the experimental data contributed the most and the results of interpreting instances of correctly and misclassified users were presented. As a result, it was shown that the Random Forest classification model can interpret the reason for the classification based on the result, and suggested that SHAP can be used as an explainable technique for the user classification using user activity sensor data.

한글키워드 : 사용자 분류, 행위기반 분류, 스마트폰 센서, 랜덤 포레스트, SHAP

keywords : user classification, activity based classification, smartphone's sensors, random forest, SHAP

* 부산대학교 IT응용공학과

접수일자: 2022.11.19. 심사완료: 2022.12.08.

† 교신저자: 김영인(email: kimyi@pusan.ac.kr)

게재확정: 2022.12.20.

1. 서론

모바일 및 웨어러블 기기의 어플리케이션에서 편리하고 지속적이면서 에너지 효율이 높은 사용자 분류 기술의 중요성이 증가하고 있다. 이를 위해 기계학습과 딥러닝 기술을 이용한 연구가 활발히 진행되고 있으며, 이를 위하여 대규모 데이터 수집과 사용자를 구분할 수 있는 분류 기술이 연구되어 성능을 개선하고 있다[1,2]. 특히 스마트폰 등 다양한 모바일과 웨어러블 기기에 탑재된 움직임 관련 센서를 사용한 연구에서 사용자의 움직임 데이터를 가지고 연속적으로 사용자를 구분할 수 있어 관련 기기를 사용하는 동안 기존 방법보다 편리하게 분류할 수 있어 더 나은 장점을 가지고 있다[3].

사용자의 움직임을 측정하여 사용자를 분류하는 기술은 생체인증(biometric) 기술에 포함되며 행동 생체 인식 기반의 지속적인 인증(continuous authentication) 기술 중 하나로 합법적이지 않은 사용자를 찾기 위하여 비정상 패턴을 식별하는 기술과 사용자 분류 기술의 두 가지가 주로 연구되고 있으며, 이 기술은 안전하고, 지속적이면서 투명하고 효율적이어야 한다[3]. 이를 위해서는 사용자간의 유사성과 차이점을 분석하여 유사한 특성을 갖고 있어 잘못 분류되는 사용자들에 대한 분석이 중요하다. 따라서 이와 같은 유사한 사용자를 파악하기 위해서는 분류 과정과 결과를 설명할 수 있어야 한다. 또한, 사용하는 기기가 소형 모바일 또는 사물인터넷 기기이므로 성능은 어느 정도 우수하면서 제한된 컴퓨팅 및 에너지 자원을 가지고 처리할 수 있는 기법이어야 한다. 그러므로 딥러닝 등 요구하는 기억 장소와 계산량이 상대적으로 큰 기법보다는 보다 적게 사용하면서 좋은 성능을 제공할 수 있는 기법에 대한 연구가 필요하다[4,5].

지금까지 사용자 분류 기술의 연구는 정확도

를 향상시키기 위하여 계산 능력이 우수한 컴퓨팅 환경을 사용한 연구가 주로 진행되고 있으며 [3,6], 계산량을 줄여 에너지 소모도 줄이기 위한 연구도 진행되고 있다[6]. 그러나 투명하게 사용자 분류의 과정이 어떻게 이루어지며 중요한 요소는 무엇인지 등 그 결정 과정과 결과를 이해할 수 있도록 글로벌 해석해 줄 수 있는 연구는 미흡한 실정이므로 이에 대한 연구가 필요한 실정이다.

최근들어 기계학습을 이용한 연구에서 입력 데이터로부터 어떠한 과정에 의해 결과가 도출되었는지 그 과정을 설명할 수 있는 XAI(eXplainable Artificial Intelligence) 기법에 대한 관심이 높아지고 있다[7]. XAI는 딥러닝을 포함한 기계학습 모델이 내놓은 결과를 사람이 이해할 수 있도록 해석하는데 필요한 기술을 포함하며, 이를 활용하면 기계학습 모델이 제시한 결과가 그 과정에서 어떤 부분을 중요하게 판단하였는지 알 수 있도록 보여주며 또한 어떤 부분으로 인하여 성공과 실패가 이루어졌는지를 설명할 수 있다[8]. 그러나 최근 관련 연구에서 사용된 기계학습 모델은 대부분 크고 복잡한 기법을 사용하여 지속적으로 사용하기에 한계가 있으며 [4,5], 제한한 분류 모델로부터 최종 사용자 분류 결과를 가져온 근거와 그 과정을 설명하기 어려워 사용자가 명확하게 해석하여 이해할 수 있도록 설명하는데도 한계가 있다. 특히 지속적인 사용자 분류의 주요 분야인 가속도계와 자이로스코프 센서로부터 수집한 사용자의 움직임 데이터를 가지고 처리하는데 있어서 사용하는 기기의 처리 능력보다 계산량이 많고 기억장소 요구량이 큰 기법보다는 상대적으로 적으면서 성능이 우수한 랜덤 포레스트(Random Forest) 알고리즘[9]과 같은 기법을 사용하면서 분류 과정도 이해하기 쉽도록 설명할 수 있는 연구가 필요하다.

이에 본 연구에서는 가속도계와 자이로스코프

센서 데이터를 랜덤 포레스트 알고리즘을 이용하여 사용자를 분류하는 과정을 시각화하여 그 과정을 해석가능하도록 하는 방법을 제안하고자 한다. 이를 위하여 사용자 데이터는 공개 데이터셋을 사용하며, XAI 기술 중 하나인 SHAP (SHapley Additive exPlanations)을 활용하여 스마트폰과 모바일 기기 사용 환경에서 사용자 움직임 센서로부터 수집한 데이터의 특징이 사용자 분류에 영향을 미치는 기여 정도를 구하여 성공과 실패에 영향을 미치는 특징을 해석할 수 있는 결과를 시각화하여 제시한다.

본 논문의 구성 및 내용은 다음과 같다. 2장에서는 연구 배경으로 사용자 움직임 관련 사용자 식별에 대한 관련 연구와 랜덤 포레스트 기법, 설명가능한 XAI 기법으로 SHAP 기술에 대하여 설명한다. 다음으로 3장에서는 본 논문에서 사용한 데이터와 XAI를 결합한 사용자 분류 과정을 설명한다. 4장에서는 개발한 분류 모델을 이용한 실험 결과와 SHAP을 활용한 결과를 설명하고, 마지막으로 5장에서는 결론과 향후 연구 방향을 제시한다.

2. 연구 배경

2.1 사용자 분류 관련 연구

사용자가 모바일, 웨어러블 및 사물인터넷 기기에 접근할 수 있는 권한이 있는지 여부를 지속적으로 확인하기 위하여 현재 사용자가 누구인지 분류하는 연구는 이상 탐지와 사용자 분류로 볼 수 있다[3]. 여기서 이상 탐지와 사용자 분류는 정상 사용자의 패턴과 비정상 패턴을 판별하는 것과 사용자 분류 모델이 합법적인 사용자와 아닌 경우를 구분할 수 있도록 사용자간의 차이를 학습하여 사용자 분류 모델은 다음과 같이 연구되었다[3].

걸음 걸이를 포함한 사용자의 움직임을 기반으로 한 연구는 움직임 패턴을 분석하여 사용자를 분류하는 기술로 데이터 수집 방법에 따라 카메라, 바닥 센서 및 웨어러블 기기를 사용하는 방법으로 분류할 수 있으며, 이 중 카메라와 바닥 센서 기반 방법은 환경이 제한되며 높은 성능의 기기를 요구하여 사용에 한계가 있다[3,10]. 그러나 모바일 기기의 센서를 활용한 연구는 사용자가 움직일 때 발생하는 신호를 신체에 부착된 움직임 관련 센서를 이용하여 사용자를 지속적으로 분류하므로 상대적으로 환경 등의 제한이 적어 이에 대한 연구가 활발히 수행되고 있으며, 이와 관련하여 특히 스마트폰을 이용한 연구로 스마트폰에 탑재된 움직임 관련 센서를 사용하므로 부피가 작고 저렴하며, 기존 스마트폰을 이용할 수 있는 장점을 가지고 있다[3,11,12,13]. 그러나 스마트폰, 웨어러블 및 IoT 기기는 메모리와 배터리 용량 및 계산 능력에 한계가 있어 현재 우수한 성능을 보인 딥러닝 모델은 자원이 한정된 모바일 및 웨어러블 기기에서는 작동하기 어려운 한계를 가지고 있다[4,5,14]. 또한 이와 같은 연구에서 개발한 사용자 분류 모델이 어떠한 과정을 수행하며, 그 과정에서 중요하게 다루어진 특징과 분류 결과에 대한 이해에 필요한 부분을 설명할 수 있도록 시각화하여 제공하지 않았다. 따라서 본 연구에서는 기존 연구에서 딥러닝으로 개발한 분류 모델을 제외한 보다 단순하면서 우수한 성능을 보인 랜덤 포레스트 기법을 사용하여 사용자 분류 과정 및 결과를 시각화하여 설명 가능하도록 하고자 한다.

2.2 랜덤 포레스트

랜덤 포레스트 알고리즘은 기계학습의 앙상블 기반 기법으로 기존 의사결정트리의 문제점인 과적합(overfitting)과 정확도 저하 등의 한계를 극복하기 위하여 개발되었다[9]. 이를 위하여 랜덤

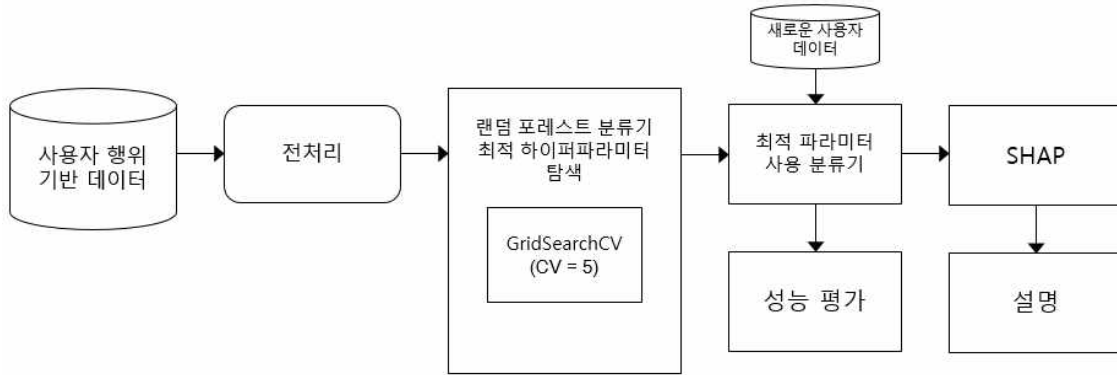


그림 1. 사용자 분류 및 해석 단계

Fig. 1. User classification and interpretation phases

포레스트 알고리즘은 입력 데이터로부터 무작위 추출한 데이터를 다수의 의사결정트리 중 무작위로 선정하여 나누어 처리하는 배깅(bagging)을 수행하며, 마지막 단계에서 투표와 같은 방식으로 다수의 의사결정트리에서 도출된 결과를 종합하여 사용한다. 이러한 과정에서 학습 데이터를 무작위로 활용하여 의사결정트리가 데이터를 분류할 수 있도록 크기를 증대시켜 과적합의 위험을 가지고 있지만, 의사결정트리의 개수를 많이 사용하고 다수결로 가장 많은 결과를 선택하여 기존 의사결정트리가 가진 학습 데이터에 과대하게 적합한 결과를 도출하여 새로운 데이터에 대해 성능이 저하되는 과적합 문제를 개선한다 [15,16].

2.3 SHAP

SHAP은 기계학습 모델의 출력을 투명하게 해석하는 XAI 방법 중 하나로 기존 게임이론에서 사용하고 있는 방법으로 각 특징의 중요도를 알기 위하여 해당 특징의 유무에 따른 기계학습 모델의 변화를 구하여 얻어낸 결과에 대한 각 특징의 기여도 값인 샐플리 값(shapley value)를 이용한 알고리즘이다[17]. 샐플리 값은 각 특징을

포함하거나 제외할 수 있는 경우의 수 만큼 계산하여 구하며 이를 통해 입력된 각 특징 변수가 개발한 모델의 결과에 어느 정도로 영향을 미쳤는지 설명한다[18]. 본 논문에서 사용하는 랜덤 포레스트 알고리즘은 트리 기반 모델이므로 이에 적합하도록 SHAP을 확장한 TreeSHAP[19,20]을 사용하여 분류 모델을 해석하여 설명한다. TreeSHAP은 트리 기반의 모델에 대하여 각 특징 변수의 기여도를 설명하는 각 인스턴스의 샐플리 값을 구하는 방법으로 기존 KernelSHAP보다 트리 기반 기계학습 알고리즘에서 보다 빠르며 정확하게 계산할 수 있다[19,20].

3. 연구 방법

3장에서는 본 논문에서 제안하는 사용자 분류 모델의 성능을 측정하고 개발한 모델을 설명할 수 있도록 만들기 위하여 그림 1과 같은 단계로 구성한다.

3.1 데이터

사용자 움직임 데이터는 사용자 움직임 인식

(HAR: Human Activity Recognition) 연구에 널리 사용되고 있는 UCI Machine Learning Repository의 공개 데이터셋인 [21]을 사용한다. 이 데이터에는 19세에서 48세 사이의 30명의 사용자가 삼성 갤럭시 S2 스마트폰을 가지고 걷기, 계단 내려오기, 계단 오르기, 눕기, 서기와 앉기를 동작하면서 3축 가속도와 자이로스코프 센서의 데이터를 수집하여 특징을 추출한 데이터이다. 사용한 센서의 신호는 50Hz로 샘플링한 후, 샘플 윈도우는 약 2.56초 단위로 50%가 겹치도록 하여 수집하였다. 가속도 센서로부터 얻은 데이터는 Acc로 그리고 자이로스코프 센서로부터 얻은 데이터는 Gyro로 표기되어 있다. 데이터셋에 있는 주요 특징의 개수는 561이며 평균, 표준편차, 엔트로피, 최솟값, 최댓값, 가중평균과 각 윈도우의 신호를 평균화하여 구한 angle 등으로 시간 영역으로부터 구한 특징은 t로 시작하는 이름을 사용하고, 각도(angle) 값을 기반으로 하는 특징은 a로 그리고 주파수 영역 기반의 특징은 f로 시작하는 특징 이름을 사용하고 있다. 본 논문은 사용자 분류 과정과 결과에서 특징의 의미를 설명하기 위한 연구이므로 561개의 특징을 모두 사용한다.

3.2 전처리

3.1절에서 설명한 바와 같이 [21]의 데이터는 여섯 가지 행동을 인식하기 위하여 수집한 데이터이므로, 본 연구의 목적인 사용자 분류에 적합하도록 레이블을 사용자 식별 번호로 변경한다. 또한 기존에 사용자 움직임 인식을 위하여 분리되어 있는 학습용과 테스트용 데이터 셋을 하나로 합쳐 사용자 분류에 적합하도록 한다. 표 1은 사용자 분류용으로 전처리를 마친 데이터셋의 각 사용자에 대한 데이터 개수를 나타낸 것으로, 전체 데이터 개수는 10,299이며, 사용자당 평균 데이터 개수는 343.3이고 표준편차는 35.71이다.

표 1. 각 사용자 데이터 전체 개수
Table 1. Total Counts of Each UserID

UserID	Total Count	UserID	Total Count	UserID	Total Count
U01	347	U11	316	U21	408
U02	302	U12	320	U22	321
U03	341	U13	327	U23	372
U04	317	U14	323	U24	381
U05	302	U15	328	U25	409
U06	325	U16	366	U26	392
U07	308	U17	368	U27	376
U08	281	U18	364	U28	382
U09	288	U19	360	U29	344
U10	294	U20	354	U30	383

3.3 실험 분류 방법 및 SHAP

랜덤 포레스트 알고리즘은 부트스트랩, 트리 최대 깊이, 최소 샘플 분할값, 트리 개수 등 하이퍼파라미터 설정이 중요하며 특히 과적합을 방지하도록 하여야 한다. 본 논문에서는 사용자 분류 모델을 최적화하기 위하여 하이퍼파라미터 튜닝 방법으로 GridSearchCV를 사용한다. 또한 학습과 테스트 과정에서 일부 데이터에 편중된 결과를 방지하기 위하여 5겹 교차검증을 사용한다. 최적 하이퍼파라미터를 구하는 과정에서 성능 평가 지표는 정확도(accuracy)를 사용한다. 다음으로 구해진 최적 하이퍼파라미터를 적용한 랜덤 포레스트 분류 모델의 사용자 분류 실험을 수행하며 그 결과인 성능 평가 지표로는 정화도 이외에 재현율(recall), 정밀도(precision), F1-점수, ROC-AUC를 사용하여 개발한 분류 모델의 성능 수준을 제시한다. 최적 하이퍼파라미터를 사용한 랜덤 포레스트 분류 모델에 XAI 기술인 SHAP을 적용하여 사용자 분류 결과에 대한 상세한 내용을 해석할 수 있도록 그림 형태로 시각화하여 제시한다. 이를 활용하여 올바른 분류와 잘못된 분류의 결과만이 아니라 다양한 시각적 형태를 보고 어떤 부분이 중요한 영향을 주었으며 그 결

과를 다양하게 해석할 수 있도록 한다.

4. 실험 및 결과

4.1 사용자 분류 모델 개발

실험 환경으로 윈도우즈 11에 파이썬 3.7 그리고 Scikit-Learn 1.0 버전[22]을 사용하였다. 먼저 성능이 우수한 랜덤 포레스트 분류 모델의 하이퍼파라미터를 찾기 위하여 Scikit-Learn의 GridSearchCV의 하이퍼파라미터 탐색은 `n_estimators`를 [10, 50, 100, 200, 300, 400, 500, 600, 700, 800], `max_features`는 [gini, sqrt], `max_depth`는 [2, 3, 4], `min_samples_split`는 [2, 3], `min_samples_leaf`는 [1, 2, 3], `bootstrap`은 [False, True], 그리고 5겹 교차검증 방법을 이용하였다. 실험한 결과 찾아낸 하이퍼파라미터는 `bootstrap`은 false, `max_features`는 sqrt, `min_samples_leaf`는 1, `min_samples_split`는 2이었으며, `n_estimators`는 400 이상에서 성능 향상이 정체되어 400을 선정하였다. 나머지 하이퍼파라미터는 기본 값을 사용하였다.

4.2 분류 결과

구해진 하이퍼파라미터를 적용하여 사용자 분류 모델을 개발하여 사용자 분류 실험을 진행하였다. 이 실험은 SHAP을 이용[23]하여 사용자 분류 과정 및 결과에 대한 설명을 시각화하여 제공하기 위한 것으로 한가지 분류 모델을 개발하여 수행하였으며, 이를 위하여 훈련과 테스트 데이터를 무작위로 75%와 25% 비율로 나누어 진행하였다. 실험 결과, 사용자 분류 모델의 성능은 정확도 90%, 재현율 0.9, 정밀도 0.9, F1-점수 0.9 및 one-vs-rest ROC의 AUC 점수는 0.994345로 비교적 높은 성능을 나타내었다.

4.3 분류 모델 및 결과 해석

랜덤 포레스트 기반 분류 모델에 어떤 특징의 기여도가 결과에 중대한 영향을 미쳤는지 해석하는 것은 중요하다. 따라서 사용자 분류 과정에서 각 특징의 기여도를 반영한 새폴리 값을 계산하여야 하며, 그림 2는 글로벌 해석을 위하여 구해진 결과를 나타낸 것이다. 그래프의 x축은 SHAP 값의 절대값 합계에 대한 평균을 나타내어 분류 결과에 기여한 특징의 정도를 바 형태로 나타낸다. 각 바 그래프에 나타난 색상은 각 사용자에게 대한 기여정도를 표시하며, 사용자 분류 모델의 기여도는 `tGravityAcc-max()-Z`, `tGravityAcc-mean()-Z`, `angle(Z, gravityMean)`, `tGravityAcc-min()-Z` 순서로 높았다. 움직임 센서의 데이터 중 Z축과 연관된 특징의 기여도가 크며, 다음으로 Y축과 관련한 특징이 중요한 것으로 나타났다. 또한 Acc로 표기된 가속도 센서로부터 추출한 특징이 상위에 위치하고 있어, 자이로스코프 센서보다 기여가 큰 것을 보여주었다.

다음으로 특징의 종속성을 분석하였다. 즉, 사용자 분류에 사용된 특징간의 상호 작용을 파악하기 위한 글로벌 해석으로 그 결과는 그림 3과 같다. 이 분석은 전역 방법으로 모든 사용자의 데이터를 가지고 분류 결과와 특징의 전역 관계에 대한 설명을 제공하며, 그림에 나타난 붉은 색과 파란 색의 점들은 `tBodyAcc-mean()-X`와 `tBodyAcc-mean()-X`의 SHAP 값을 보이며, `tBodyAcc-mean()-X`의 기여도가 0.10 부근에서 기여도가 증가하여 0.65에서 감소하는 것을 확인할 수 있었으며, `tBodyGyroJerkMag-entropy()` 사이의 0.25부근에서 파란색의 뚜렷한 세로 패턴이 생긴 것으로 명확한 상호작용이 있었음을 알 수 있었다.

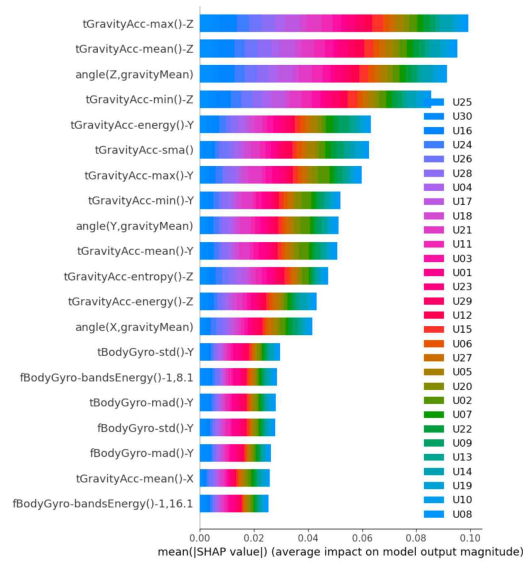


그림 2. 주요 특징의 기여도
Fig. 2. Contribution of important features

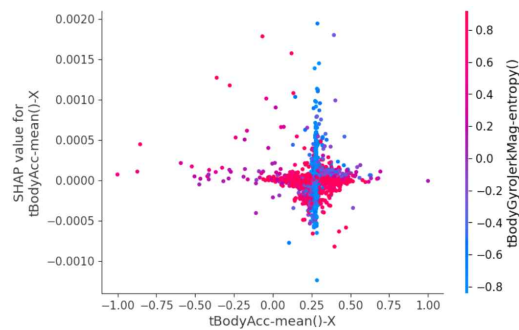


그림 3. 특징의 종속성
Fig. 3. Dependencies of features

각 사용자를 구분하는데 기여도가 높은 각 특징에 대한 요약의 예는 그림 4와 같다. 이 그림에 있는 각 점은 각 특징에 대한 새플리 값으로 x축은 새플리 값 그리고 y축은 특징 값에 의해 결정되며, 각 색상은 특징 값의 높고 낮음을 표시하며 특징의 기여 순서에 따라 나타난다. 이 해석을 통하여 각 특징이 얼마나 결과에 영향을 주었는지를 알아볼 수 있다. 이 그림을 통해 이 예제의 사용자에게 대한 분류가 tGravityAcc-

max()-Y, tGravityAcc-mean()-Z, angle(Z, gravityMean) 순서로 중요했음을 알 수 있었으며, 이러한 설명 과정을 통하여 사용자 단위로 특징의 기여 정도를 구하여 설명할 수 있다.

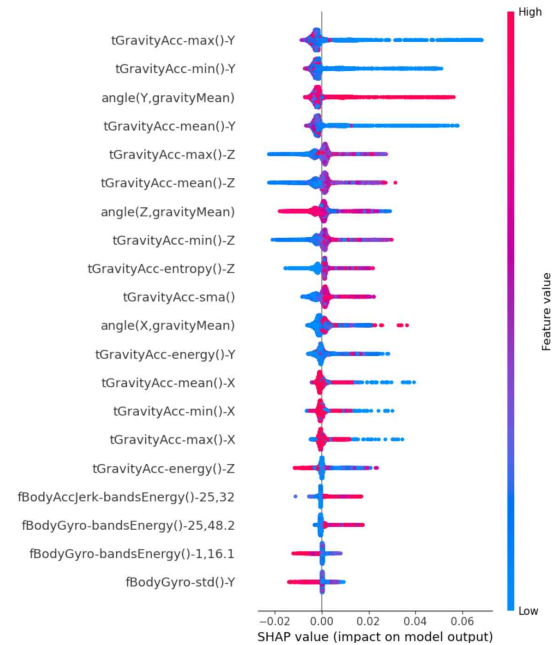


그림 4. 사용자별 주요 특징의 기여도
Fig. 4. Contribution of important features by User

그림 5는 로컬 해석으로 테스트 데이터에 대한 사용자 분류 결과에서 가장 성능이 낮았던 8번 사용자의 데이터 중 잘못 분류한 인스턴스에 대한 그림 5의 (a)와 반대로 분류가 잘 되었던 21번 사용자의 인스턴스를 나타낸 것은 (b)이다. 그림에서 굵은 글씨의 숫자는 사용자 분류의 기본 값으로 붉은색 화살표는 해당 사용자에게 대한 분류 가능성을 증가시키는 특징의 영향 정도를 나타내고, 파란색 화살표는 사용자 분류 가능성을 줄이는 특징의 영향을 나타낸다. 각 화살표의 크기는 특징이 차지하는 상대적인 크기로 해당 특징의 영향 정도를 나타낸다. 여기서 그림 5의 (a)는 기준값 0.91보다 해당 특징이 높은 영향을 주

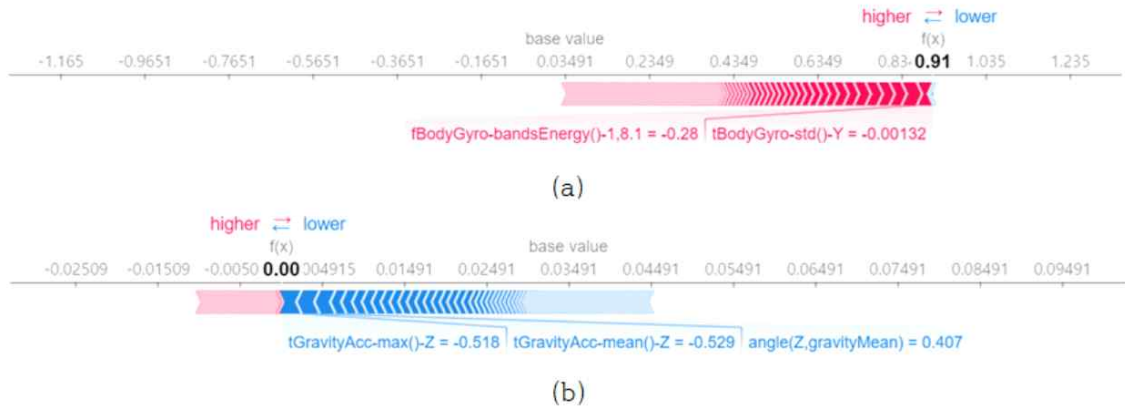


그림 5. 테스트 데이터의 사용자 분류 결과 예
Fig. 5. Example results of user classification in test data

어 잘못 분류되었음을 나타내며, (b)는 0.0의 기준값보다 높은 해당 특징의 순서로 사용자 분류 결과에 큰 영향을 주는 특징으로 해석하여 설명할 수 있다.

마지막으로 그림 6은 그림 5에 대한 다른 표현으로 SHAP의 폭포수 플롯이다. 그림 5와 같이 (a)는 분류가 잘 안되었던 사용자 8번을 나타내고, 사용자의 인스턴스 및 분류가 잘 되었던 21번 사용자는 (b)로 나타내었다. 그림과 같이 y축에 있는 특징들이 긍정적으로 영향을 준 경우에는 붉은색으로 나타내고 반대로 부정적인 기여를 한 경우에는 파란색으로 나타내며, 그 값을 화살표 옆에 표시한다. 따라서 이 그래프를 통해 관심이 있는 사용자의 분류가 어떤 특징에 영향을 받아 이루어졌는지를 해석하여 설명할 수 있다.

4. 결론

본 논문에서는 사용자 움직임을 가지고 지속적으로 사용자를 식별하기 위한 연구로, 스마트폰에 탑재된 가속도와 자이로스코프 센서로부터 수집한 데이터를 가지고 사용자 분류 모델을 개

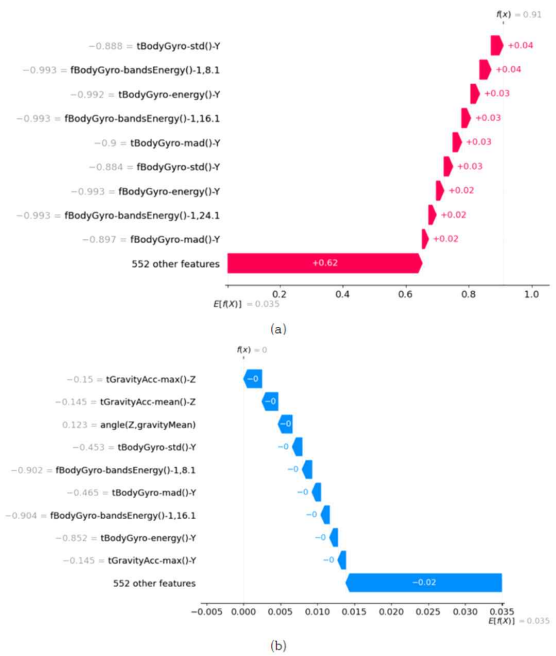


그림 6. 테스트 데이터 사용자별 폭포수 형태 예
Fig. 6. Examples of Waterfall plot for each user in the test data

발하고 사용자를 분류하는 과정을 해석하여 설명하는데 도움을 줄 수 있는 XAI 기법인 SHAP을 활용한 방법을 제시하였다. 실험을 통하여 분류

모델과 해석 과정을 제시하여 사용자 분류 과정에서 각 특징이 기여한 정도를 보여 분류 결과에 어떤 영향을 미쳤는지를 설명할 수 있음을 보였다. 이 결과를 이용하여 다른 사용자가 같은 사용자로 분류된 경우에 대해 그 원인을 분석하여 활용한다면 사용자를 구분하는 기술의 발전에 기여할 수 있을 것으로 기대한다. 또한 제안한 방법으로 좀 더 다양하고 많은 사용자의 움직임 데이터를 수집하여 실험한다면 지금보다 그 과정을 투명하게 해석하며 분석할 수 있어 신뢰도가 개선된 사용자 분류 기법 개발에 도움을 줄 것으로도 기대한다.

이 논문은 부산대학교
기본연구지원사업(2년)에 의하여 연구되었음.

참 고 문 헌

- [1] M. Abuhamad, A. Abusnaina, D. Nyang & D. Mohaisen, (2020). Sensor-based continuous authentication of smartphones' users using behavioral biometrics: A contemporary survey, *IEEE Internet of Things Journal*, 8(1), 65-84. DOI: 10.1109/JIOT.2020.3020076
- [2] Ahmed Fraz Baig & Sigurd Eskeland, (2021). Security, privacy, and usability in continuous authentication: a survey, *Sensors*, 21(17). DOI:10.3390/s21175967
- [3] Y. Liang, S. Samtani, B. Guo & Z. Yu, Behavioral Biometrics for Continuous Authentication in the Internet-of-Things Era: An Artificial Intelligence Perspective, *IEEE Internet of Things Journal*, 7(9). 9128-9143. DOI: 10.1109/JIOT.2020.3004077
- [4] M. EhatishamulHaq, M. A. Azam, J. Loo, K. Shuang, S. Islam, U. Naeem, Y. Amin, Authentication of smartphone users based on activity recognition and mobile sensing, *Sensors*, Vol. 17, No. 9: 2043, 2017. DOI:10.3390/s17092043
- [5] K. Chen, D. Zhang, L. Yao, B. Guo, Z. Yu & Y. Liu, (2020). Deep learning for sensor-based human activity recognition: overview, challenges and opportunities, *arXiv preprint arXiv:2001.07416*. DOI:10.1145/1122445.1122456
- [6] A. Rawal, J. Mccoy, D. B. Rawat, B. Sadler & R. Amant, (2021). Recent advances in trustworthy explainable artificial intelligence: Status, challenges and perspectives, *IEEE Trans. Artif. Intell.*, No. 4. DOI:10.1109/TAI.2021.3133846
- [7] I. Kok, F. Y. Okay, O. Muyanli & S. Ozdemir, (2022). Explainable artificial intelligence (xai) for internet of things: a survey, *arXiv preprint arXiv:2206.04800*. DOI:10.48550/arXiv.2206.04800
- [8] Alejandro Barredo Arrieta, Natalia Díaz-Rodríguez, Javier Del Ser, Adrien Bannetot, Siham Tabik, Alberto Barbado, Salvador Garcia, Sergio Gil-Lopez, Daniel Molina, Richard Benjamins, Raja Chatila & Francisco Herrera, (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI, *Information Fusion* Vol. 58, 82-115. DOI: 10.1016/j.inffus.2019.12.012
- [9] L. Breiman, (2001). Random forests, *Machine Learning*, 45(1), 5-32. DOI:10.1023/A:1010933404324
- [10] Z. Wu, Y. Huang, L. Wang, X. Wang & T. Tan, (2017). A comprehensive study on cross-view gait based human identification with deep CNNs, *IEEE Trans. Pattern Anal. Mach. Intell.*, 39(2), 209-226. DOI: 10.1109/TPAMI.2016.2545669
- [11] W. Meng, D. S. Wong, S. Furnell & J. Zhou, (2014). Surveying the development of biometric user authentication on mobile phones, *IEEE Commun. Surveys Tuts.*,

- 17(3), 1268 - 1293. DOI: 10.1109/COMST.2014.2386915
- [12] M. D. Marsico, A & Mecca, (2019). A survey on gait recognition via wearable sensors, *ACM Comput. Surveys*, 52(4), 1 - 39. DOI:10.1145/3340293
- [13] M. Muaaz & R. Mayrhofer, (2017). Smartphone-based gait recognition: From authentication to imitation, *IEEE Trans. Mobile Comput.*, 16(11), 3209 - 3221. DOI: 10.1109/TMC.2017.2686855
- [14] N. D. Lane, S. Bhattacharya, P. Georgiev, C. Forlivesi & F. Kawsar, (2015). An early resource characterization of deep learning on wearables, smartphones and Internet-of-Things devices, *Proc. Int. Workshop Internet Things Towards Appl. (IoT-App)*, 7 - 12. DOI: 10.1145/2820975.2820980
- [15] T. K. Ho, (1995). Random Decision Forests, 3rd International Conference on Document Analysis and Recognition, 278 - 282. DOI: 10.1109/ICDAR.1995.598994
- [16] R.-C. Chen, C. Dewi, S.-W. Huang & R. E. Caraka, (2020). Selecting critical features for data classification based on machine learning methods, *Journal of Big Data*, 7(1). DOI: 10.1186/s40537-020-00327-4
- [17] S. M. Lundberg & S. I. Lee, (2017). A unified approach to interpreting model predictions, *Advances in neural information processing systems*, 30. DOI: 10.48550/arXiv.1705.07874
- [18] M Lundberg Scott, Gabriel Erion, Hugh Chen, Alex DeGrave, Jordan M Prutkin, Bala Nair, Ronit Katz, Jonathan Himmelfarb, Nisha Bansal & Su-In Lee, (2019). Explainable ai for trees: From local explanations to global understanding, *arXiv preprint arXiv:1905.04610*. DOI: 10.48550/arXiv.1905.04610
- [19] Scott M Lundberg, Gabriel G Erion & Su-In Lee. (2019). Consistent individualized feature attribution for tree ensembles. *arXiv preprint arXiv:1802.03888*. DOI:10.48550/arXiv.1802.03888
- [20] Scott M Lundberg, Gabriel Erion, Hugh Chen, Alex DeGrave, Jordan M Prutkin, Bala Nair, Ronit Katz, Jonathan Himmelfarb, Nisha Bansal & Su-In Lee, (2019). Explainable ai for trees: From local explanations to global understanding, *arXiv preprint arXiv:1905.04610*. DOI:10.48550/arXiv.1905.04610
- [21] D. Anguita, A. Ghio, L. Oneto, X. Parra & J. L. Reyes-Ortiz, (2013). A public domain dataset for human activity recognition using smartphones, *European Symposium on Artificial Neural Networks*, 437-442. DOI:10.3390/s20082200
- [22] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel & E. Duchesnay, (2011). Scikit-learn: Machine learning in Python, the *Journal of machine Learning research*, 12, .2825-2830. DOI:10.48550/arXiv.1201.0490
- [23] Idit Cohen, (2021). Explainable AI (XAI) with SHAP -Multi-Class Classification Problem, <https://towardsdatascience.com/explainable-ai-xai-with-shap-multi-class-classification-problem-64dd30f97cea>.

저 자 소 개



김영인(Young-In Kim)

1996.2 명지대학교 컴퓨터공학과 박사
2007.8-2008.7 Univ. of Missouri 방문교수
2006.3-현재 : 부산대학교 교수
<주관심분야> 데이터베이스, 데이터마이닝