

임계값 설정을 통한 근치적 위절제술 후 합병증 발생 예측 모델의 성능 평가

(Performance of a Model to Predict Complication Occurrence after Radical Gastrectomy according to Thresholds)

임수연*, 최자윤**

(Suyeon Lim, Ja Yun Choi)

요약

근치적 위절제술 후 합병증은 환자의 회복을 방해하고 예후를 악화시키므로, 이를 사전에 예측하는 것이 중요하다. 본 연구는 4,892명의 전자의무기록 데이터를 활용하여 근치적 위절제술 후 합병증을 예측하는 랜덤포레스트 모델을 개발하고, 데이터 불균형 문제를 해결하기 위해 임계값을 설정하여 성능을 개선하였다. 연구 결과, 근치적 위절제술 후 합병증 발생률은 17.2%였고 최적의 임계값인 0.31에서 F1 score, 정밀도, 재현율이 향상되었다. 본 연구에서 임계값 설정은 의료 데이터와 같이 불균형이 심한 데이터에 활용되어 예측 모델의 성능을 향상시켰다. 향후 임계값 설정과 더불어 다양한 불균형 처리 기법의 성능 비교가 필요하다.

■ 중심어 : 데이터 불균형 ; 임계값 ; 머신러닝 모델

Abstract

Postoperative complications after radical gastrectomy hinder patient recovery and worsen prognosis, making early prediction critically important. This study developed a Random Forest model to predict postoperative complications using electronic medical records from 4,892 patients. To address the issue of data imbalance, we adjusted the classification threshold to improve model performance. The results showed that postoperative complications after radical gastrectomy occurred in 17.2% of patients, and that the optimal threshold at 0.31 improved F1 score, precision, and recall, thereby enhancing prediction performance. This study suggests that threshold adjustment plays a crucial role in improving prediction models from medical data with imbalance. Further research is needed to compare the classification threshold with various imbalance handling techniques.

■ keywords : Data Imbalance ; Threshold ; Machine Learning Model

1. 서론

근치적 위절제술은 위암 환자에게 적용되는 표준적인 치료 방법으로, 병변의 완전한 절제와 장기 생존을 목표로 한다[1]. 그러나 수술 후에는 다양한 합병증이 발생할 수 있으며, 이는 환자의 회복을 지연시키고 전반적인 예후를 악화시키는 주요 요인이다[2]. 따라서 수술

후 합병증의 발생 가능성을 사전에 예측하는 것은 환자의 치료 결과를 향상시키며, 치료 계획을 최적화하는 데 있어 매우 중요하다[3].

그동안 수술 후 합병증을 예측하기 위해 여러 방법이 시도되었지만, 기존 방법들은 의료 데이터에 내재된 복잡하고 비선형적인 변수 간의 관계를 충분히 반영하지 못하는 한계가 있다[4]. 이를 보완하기 위해 머신러닝 기법

* 정회원, 미시간대학교 간호대학 박사후연구원

** 정회원, 전남대학교 간호대학 교수

이 논문은 2024년 박사학위논문의 일부를 수정·보완한 것임.

이 논문은 2024년 한국정보처리학회 학술대회에서 발표한 논문을 수정·보완한 것임.

접수일자 : 2025년 04월 25일

수정일자 : 2025년 05월 27일

게재확정일 : 2025년 06월 23일

교신저자 : 최자윤 e-mail : choijy@jnu.ac.kr

이 주목받고 있으며, 그 중에서도 랜덤포레스트는 높은 예측 정확도와 안정성을 바탕으로 의료 분야에서 널리 활용되고 있다[5,6]. 랜덤포레스트는 여러 의사결정 트리를 결합해 예측 성능을 향상시키며, 무작위로 추출된 데이터 샘플과 변수 조합을 통해 과적합을 방지하고 모델의 일반화 가능성을 높인다[7]. 또한 의료 데이터 내 포함되는 노이즈와 고차원 데이터, 비선형적 특성에도 견고한 성능을 보인다[6]. 아울러 변수 중요도를 정량적으로 산출할 수 있어, 의료진의 의사결정 지원에도 유용하다[5,7]. 이러한 특성을 바탕으로 랜덤포레스트는 복잡하고 가변성을 지닌 의료 데이터를 기반으로 한 예측 모델 개발에 적합한 알고리즘으로 생각된다.

한편, 의료 데이터는 정상 사례가 대부분을 차지하고 질환 사례는 상대적으로 적게 나타나기 때문에 데이터 불균형 문제가 자주 발생한다[8]. 이로 인해 모델은 다수 클래스를 중심으로 학습하게 되어, 실제로 중요한 소수 클래스인 질환 사례를 제대로 예측하지 못하는 문제가 발생할 수 있다[9]. 이러한 문제는 환자의 생명과 직결될 수 있기 때문에, 의료 분야에서는 데이터 불균형 문제를 효과적으로 해결하는 것이 필수적이다[10]. 이를 해결하기 위해 리샘플링(resampling) 기법, 비용 민감 학습(cost-sensitive learning), 앙상블 학습(ensemble learning) 등의 방법이 제안되어 왔다[11]. 그러나 이러한 방법들은 대부분 모델 학습 이전 단계에서 데이터나 가중치를 조정하는 방식으로 적용되기 때문에, 데이터 특성의 왜곡, 모델 복잡도 증가, 해석의 어려움 등의 한계를 가진다[11].

반면, 임계값 설정은 학습된 모델의 예측 확률을 기반으로 분류 기준을 조정하는 후처리 방식으로, 데이터나 모델을 재구성하지 않고도 적용 가능한 효율적인 방법이다[11-13]. 특히, 데이터로부터 최적의 임계값을 도출하여

적용함으로써 예측 성능을 향상시키고, 임상적으로 중요한 소수 클래스에 대한 민감도를 높일 수 있음이 보고되었다[13-15]. 또한, 리샘플링이나 비용 민감 학습, 앙상블 학습과 비교했을 때, 그 성능이 동등하거나 더 우수하였다[14,16]. 이러한 점에서 임계값 설정은 데이터 변형으로 인한 오분류의 임상적 위험을 최소화하고, 의료 분야에서 실용적이고 적절한 접근이라 생각된다.

따라서 본 연구에서는 전자의무기록을 활용하여 근치적 위절제술 후 합병증을 예측하는 랜덤포레스트 모델을 개발하고, 데이터 불균형 문제를 해결하기 위해 적절한 임계값을 설정하여 모델의 성능을 비교 및 평가하고자 한다.

II. 관련 연구

데이터 불균형 문제를 해결하기 위한 주요 기법으로는 리샘플링 기법, 비용 민감 학습, 앙상블 학습, 임계값 설정이 있다.

1. 리샘플링 기법

리샘플링(resampling) 기법은 데이터의 불균형을 해결하기 위해 샘플 수를 조정하는 방법으로, 오버샘플링(oversampling)과 언더샘플링(undersampling)으로 구분된다[17]. 오버샘플링은 소수 클래스의 샘플을 복제하거나 새로운 샘플을 생성하여 샘플 수를 증가시키는 방법이다. 대표적인 방법으로 Synthetic Minority Oversampling Technique (SMOTE)가 있으며, 기존 데이터 포인트 사이에 새로운 샘플을 생성하는 방식으로 소수 클래스의 분포를 확장한다[18]. 반대로 언더샘플링은 다수 클래스의 샘플 수를 줄여 클래스 간 균형을 맞추는 방법이다. 주요 기법으로는 Tomek-link 방법[19]과 클러스터 중심 기법[20]이 있다. 리

샘플링 기법은 구현이 간단하고 모델 성능을 향상시키는데 유용하지만, 데이터 특성을 왜곡하거나 중요한 정보를 손실할 위험이 있어 신중한 적용이 필요하다.

2. 비용 민감 학습

비용 민감 학습(cost-sensitive learning)은 소수 클래스의 오분류 비용을 고려하여 중요한 오류에 더 큰 가중치를 부여한다. 이 방법은 단순히 데이터 분포를 조정하는 것이 아니라, 알고리즘을 수정하거나 데이터 샘플에 비용 가중치를 적용하여 모델이 특정 오류를 민감하게 인식하도록 한다[11]. 비용 민감 학습은 내부적으로 비용에 민감한 손실 함수를 알고리즘에 통합하거나, 외부적으로 오분류 비용을 반영하기 위해 훈련 데이터에 가중치를 부여하는 두 가지 방법으로 구현할 수 있다[21]. 이러한 방법은 기존의 데이터 특성을 변경하지 않고도 모델의 성능을 개선할 수 있다는 장점이 있다. 하지만 이 방법을 효과적으로 적용하기 위해서는 도메인 지식을 기반으로 오분류 비용을 적절하게 설정하는 것이 중요하다. 또한, 비용 설정과 알고리즘 수정 과정에서 모델의 복잡도가 증가하고, 특정 예측 결과가 비용 가중치에 의해 어떻게 조정되었는지 파악하기 어려워 결과의 해석에 제약이 있을 수 있다[11].

3. 앙상블 학습

앙상블 학습(ensemble learning)은 여러 분류기를 결합하여 더욱 강력한 분류기를 만들어 예측 성능을 향상시키는 방법으로, 특히 단일 분류기로는 해결하기 어려운 데이터 불균형을 처리하는데 유용하다[11]. 대표적인 기법으로는 배깅(Bagging), 부스팅(Boosting), 아다부스트(AdaBoost)가 있다.

배깅은 여러 부트스트랩 샘플을 사용하여 각기 다른 분류기를 학습시키고, 이들의 예측을 평균 또는 다수결 방식으로 결합함으로써 예측의 안정성과 정확도를 높인다. 또한, 소수 클래스와 다수 클래스의 데이터 분포를 균형 있게 사용함으로써 데이터 불균형을 해결한다[22]. 부스팅은 성능이 낮은 약한 학습자를 순차적으로 학습시켜 강한 학습자로 전환시키는 방법으로, 반복(iteration) 단계마다 이전 모델이 잘못 분류한 샘플에 집중함으로써 전체적인 오류를 점진적으로 줄인다. 특히 불균형 데이터에서는 소수 클래스의 오분류 샘플에 반복적으로 가중치가 부여되기 때문에, 학습 과정에서 소수 클래스에 대한 민감도가 자연스럽게 향상된다[23]. 아다부스트는 부스팅의 한 일종으로 오분류된 샘플에 더 큰 가중치를 부여하여 분류가 어려운 샘플에 집중하도록 유도하는 점은 부스팅과 유사하다. 다만, 최종 예측 시 각 분류기의 성능을 반영하여 성능이 높은 분류기에 더 큰 가중치를 부여한 결과를 결합한다는 점에서 차이가 있다[24].

이처럼 앙상블 학습은 단일 분류기의 한계를 극복하고, 다양한 분류기의 강점을 조합함으로써 모델의 안정성과 일반화 성능을 향상시키는데 효과적이다. 하지만, 계산 비용이 많이 들고 불균형이 심한 경우 성능 개선에 한계가 있다[11]. 또한 다수의 분류기가 결합된 복잡한 구조로 인해 예측 결과의 도출 과정이 복잡하고, 이에 따른 해석 가능성이 낮아 실제 임상 적용에는 제약이 따를 수 있다[11].

4. 임계값 설정

모델의 예측 결과는 예측 확률(prediction probability)에 따라 결정되며 임계값은 예측 확률을 기반으로 실제 클래스에 해당하는 기준점이다[12]. 일반적으로 이진 분류에서는 0.5를 기준으로 클래스를 구분하지만, 불균형 데

이터에서는 이 기준이 적절하지 않을 수 있다[13]. 따라서 불균형 데이터의 경우, 다수 클래스에 대한 편향을 줄이기 위해 모델과 데이터의 특성에 맞는 최적의 임계값을 설정하는 것이 중요하다[15]. 임계값 설정은 데이터 처리 단계가 아닌 모델을 적합하게 개발한 후에 적용하기 때문에, 데이터 세트를 조정하거나 중요한 샘플을 버리거나 기존 학습자를 수정하지 않는다는 장점이 있다[11]. 또한 임계값 설정을 통한 성능 최적화는 모델의 재학습 없이 적용이 가능하고, 단순히 예측 확률의 기준점을 변경하는 것만으로도 구현이 가능하다[13]. 실제로 van den Goorbergh 등[14]은 리샘플링과 비용 민감 학습이 민감도 향상에는 기여하였으나, 예측 확률의 왜곡을 초래할 수 있다고 하였다. 이와 달리, 임계값 설정은 민감도를 향상시키면서도 예측 확률에는 영향을 주지 않아 보다 안정적이고 해석 가능한 접근법으로 제시되었다. 또한, Zheng 등[16]은 심혈관계 임상 데이터를 활용한 연구에서 임계값 설정을 적용한 모델이 비용 민감 학습과 앙상블 모델보다 민감도 및 모델의 전반적인 판별 성능(Area Under the Receiver Operating Characteristic Curve)에서 모두 우수한 성능을 보였음을 보고하였다. 이처럼 임계값 설정은 소수 클래스의 민감도를 효율적으로 향상시키며, 기존의 리샘플링이나 비용 민감 학습, 앙상블 학습과 비교했을 때 동등하거나 더 우수한 성능을 보이는 것으로 나타났다[13,16]. 다만 최적의 임계값을 찾기 위해서는 여러 임계값을 시도해보고 성능을 평가해야 하며, 데이터의 특성에 따라 적절한 임계값이 달라질 수 있다는 점을 유의해야 한다.

III. 실험

1. 데이터 세트 생성 및 전처리

본 연구에서는 일개 대학병원에서 위암 진단

하에 근치적 위절제술을 받은 5,085명의 전자의무기록 데이터를 활용하였다. 데이터 전처리 과정에서 결측치가 존재한 193명을 제외하였으며, 최종적으로 4,892명의 데이터를 모델 개발에 사용하였다. 결과 변수인 수술 후 합병증은 퇴원 시점에 전자의무기록지에 기록된 1등급 이상의 아코디언 중증도 분류 기준에 따라 정의하였다[25]. 최종 분석 대상 중 수술 후 합병증 발생군은 843명(17.2%), 비발생군은 4,049명(82.8%)으로 데이터 불균형을 보였다.

예측 변수는 선행연구를 통해 확인된 수술 후 합병증 영향요인 중 전자의무기록에서 수집 가능한 39개의 변수를 선정하였다. 이들 변수는 수술 전, 수술 중, 수술 후 요인으로 분류하였으며, 변수 간 상관성과 모델 성능 기여도 등을 종합적으로 고려하여 최종 32개의 변수를 입력 변수로 확정하였다. 범주형 변수는 원-핫 인코딩을 통해 처리하였고, 연속형 변수는 StandardScaler를 적용하여 표준화하였다.

2. 모델 개발

본 연구에서는 모델 개발을 위해 데이터 세트는 학습용 80%, 평가용 20%의 비율로 무작위 분할하였으며, 데이터 불균형을 고려하여 계층화된 분할을 적용하였다. 분할된 학습용 데이터는 랜덤포레스트 알고리즘에 입력하여 근치적 위절제술 후 합병증 발생 예측 모델을 개발하였다. 또한, 개발된 모델 성능의 일관성과 안정성을 평가하기 위해서 stratified 5-fold cross validation을 수행하였다.

3. 성능 평가 지표

본 연구에서 개발된 모델의 성능 평가를 위해 정밀도(Precision), 재현율(Recall), F1 score, 정확도(Accuracy), 특이도(Specificity), Area Under the Receiver Operating

Characteristic Curve (AUROC)의 측정 지표를 사용하였다.

4. 임계값 설정 및 성능 비교

본 연구에서는 근치적 위절제술 후 합병증이 발생한 소수 클래스의 예측이 중요하므로, 재현율이 중요한 성능 지표였다. 하지만 재현율을 개선할 경우, 정밀도가 낮아지는 트레이드 오프(trade-off) 현상이 발생하였다. 그러므로 정밀도와 재현율 간의 균형을 고려한 최적의 모델 성능을 도출하기 위해, 두 지표의 조화평균인 F1 score를 최적의 임계값 설정 기준으로 사용하였다. 모델의 임계값 설정에 따른 정밀도-재현율의 변화 곡선은 그림1과 같으며, 최적의 임계값은 0.31로 나타났다.

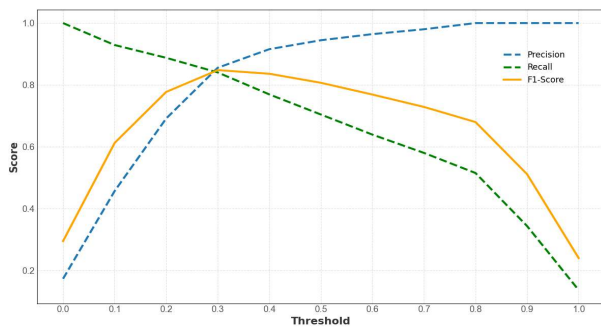


그림 1. 임계값 설정에 따른 정밀도, 재현율, F1 score의 변화

임계값 설정에 따른 모델의 성능을 평가하기 위해 최적의 임계값 0.31을 중심으로 0.11부터 0.5까지의 범위를 0.1 간격으로 변화시키며 성능을 비교하였다. 그 결과, 본 연구에서 임계값 설정에 따른 근치적 위절제술 후 합병증 발생 예측 모델의 성능 결과는 표1과 같다.

표 1. 임계값 설정에 따른 성능 비교 결과

Metrics	Threshold				
	0.11	0.21	0.31 (Best)	0.41	0.5 (Basic)
Precision	0.47	0.71	0.88	0.92	0.94
Recall	0.92	0.89	0.83	0.76	0.70
F1 score	0.62	0.79	0.86	0.83	0.81
Accuracy	0.80	0.92	0.95	0.95	0.94
Specificity	0.78	0.92	0.98	0.99	0.99
AUROC	0.95	0.95	0.95	0.95	0.95

최적의 임계값인 0.31에서 F1 score는 가장 높은 값인 0.86을 보였으며, 정밀도와 재현율은 각각 0.88과 0.83으로 나타났다. 이는 기본 임계값인 0.5를 적용했을 때의 정밀도 0.94, 재현율 0.70, F1 score 0.81로 나온 결과와 비교하면, 소수 클래스에 대한 예측 성능이 크게 향상되었음을 보여준다. 또한 임계값이 높아질수록 소수 클래스 예측에 대한 기준이 높아지면서, 재현율이 낮아짐을 확인할 수 있었다. 그러나 정확도, 특이도, AUROC와 같은 다른 성능 지표에서는 큰 차이가 없는 것으로 나타났다.

IV. 결 론

본 연구에서는 다양한 임계값 설정에 따른 근치적 위절제술 후 합병증 예측 모델의 성능을 비교 및 평가하였다. 그 결과, 최적의 임계값 설정이 전자의무기록과 같이 불균형한 데이터를 활용한 예측 모델에서 성능을 개선할 수 있음을 확인하였다. 특히, 최적의 임계값을 설정함으로써 F1 score, 재현율, 정밀도 등의 주요 성능 지표가 개선되었으며, 이는 합병증 발생이라는 소수 클래스에 대한 예측 능력이 개선되었음을 의미한다. 따라서, 임상 현장에서 보다 정확한 위험 평가 도구로 활용될 수 있음을 시사한다.

본 연구의 한계점으로는 단일 기관의 데이터만을 사용하였기 때문에, 향후 연구에서는 다기관 데이터를 활용하여 모델의 일반화 가능성을 평가하고 다양한 의료 환경에서 발생할 수 있는 데이터 불균형을 다루기 위한 보다 정교한 임계값 설정 방법을 개발할 필요가 있다. 또한, 임계값 설정 외에도 다양한 데이터 불균형 처리 기법을 적용하여 그 효과를 비교하는 연구가 필요하다. 이를 통해 의료 분야에서의 머신러닝 모델 적용 시 데이터 불균형을 해결하기 위한 방향을 제시할 수 있을 것이다.

REFERENCES

- [1] Japanese Gastric Cancer Association. "Japanese gastric cancer treatment guidelines 2021.", *Gastric Cancer*, vol. 26, no. 1, pp. 1 - 25, 2023.
- [2] S. Wang, L. Xu, Q. Wang, J. Li, B. Bai, Z. Li, X. Wu, P. Yu, X. Li and J. Yin, "Postoperative complications and prognosis after radical gastrectomy for gastric cancer: a systematic review and meta-analysis of observational studies.", *World Journal of Surgical Oncology*, vol. 17, no. 1, pp. 1 - 10, 2019.
- [3] M. Kanda, "Preoperative predictors of postoperative complications after gastric cancer resection.", *Surgery Today*, vol. 50, no. 1, pp. 3-11, 2020.
- [4] A. Géron, *Hands-on Machine Learning with Scikit-Learn, Keras, and TensorFlow*, O'Reilly Media, Inc, 2022.
- [5] A.H. Hassan, R. Sulaiman, M. Abdulhak, and H. Kahtan, "Bridging data and clinical insight: explainable AI for ICU mortality risk prediction.", *International Journal of Advanced Computer Science and Applications*, vol. 16, no. 2, pp. 743 - 750, 2025.
- [6] S. Blotwijk, C. Raets and K. Barbé, "Exploratory study on Evolutionary Random Forests for Classification in Medical Datasets.", *2023 IEEE International Symposium on Medical Measurements and Applications (MeMeA)*, pp. 1 - 6, 2023.
- [7] L. Breiman, "Random forests.", *Machine Learning*, vol. 45, pp. 5 - 32, 2001.
- [8] N. Liu, X. Li, E. Qi, M. Xu, L. Li and B. Gao, "A novel ensemble learning paradigm for medical diagnosis with imbalanced data.", *IEEE Access*, vol. 8, pp. 171263 - 171280, 2020.
- [9] M. Khalilia, S. Chakraborty and M. Popescu, "Predicting disease risks from highly imbalanced data using random forest.", *BMC Medical Informatics and Decision Making*, vol. 11, pp. 1 - 13, 2011.
- [10] D.H. Qudsi, "Predictive Analytics Data Mining in Imbalanced Medical Dataset.", *Journal Komputer Terapan*, vol. 2, no. 2, pp. 195 - 204, 2016.
- [11] H. Ali, M.M. Salleh, R. Saedudin, K. Hussain and M.F. Mushtaq, "Imbalance class problems in data mining: A review.", *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 14, no. 3, pp. 1560 - 1571, 2019.
- [12] J.L. Leevy, J.H. Johnson, J. Hancock and T.M. Khoshgoftaar, "Threshold optimization and random undersampling for imbalanced credit card data.", *Journal of Big Data*, vol. 10, no. 1, p. 58-79, 2023.
- [13] C. Esposito, G.A. Landrum, N. Schneider, N. Stiefl and S. Riniker, "GHOST: adjusting the decision threshold to handle imbalanced data in machine learning.", *Journal of Chemical Information and Modeling*, vol. 61, no. 6, pp. 2623-2640, 2021.
- [14] R. van den Goorbergh, M. van Smeden, D. Timmerman, and B. Van Calster, "The harm of class imbalance corrections for risk prediction models: illustration and simulation using logistic regression.", *Journal of the American Medical Informatics Association*, vol. 29, no. 9, pp. 1525 - 1534, 2022.
- [15] G. Mulugeta, T. Zewotir, A.S. Tegegne, L.H. Juhar, and M.B. Muleta, "Classification of imbalanced data using machine learning algorithms to predict the risk of renal graft failures in Ethiopia.", *BMC Medical Informatics and Decision Making*, vol. 23, no. 1, pp. 98 - 114, 2023.
- [16] H. Zheng, S.W.A. Sherazi, and J.Y. Lee, "A cost-sensitive deep neural network-based prediction model for the mortality in acute myocardial infarction patients with hypertension on imbalanced data.", *Frontiers in Cardiovascular Medicine*, vol. 11, pp. 1276608, 2024.
- [17] M.S. Shelke, P.R. Deshmukh and V.K. Shandilya, "A review on imbalanced data handling using undersampling and oversampling technique.", *International Journal of Recent Trends in Engineering and Research*, vol. 3, no. 4, pp. 444 - 449, 2017.
- [18] N.V. Chawla, K.W. Bowyer, L.O. Hall and W.P. Kegelmeyer, "SMOTE: synthetic minority over-sampling technique.", *Journal of Artificial Intelligence Research*, vol. 16, pp. 321 - 357, 2002.
- [19] I. Tomek, "A generalization of the k-NN rule.", *IEEE Transactions on Systems, Man, and Cybernetics*, vol. 6, no. 2, pp. 121 - 126, 1976.
- [20] G. Lemaître, F. Nogueira and C.K. Aridas, "Imbalanced-learn: A python toolbox to tackle the curse of imbalanced datasets in machine learning.", *Journal of Machine Learning Research*, vol. 18, no. 17, pp. 1 - 5, 2017.
- [21] B. Zadrozny and C. Elkan, "Learning and

making decisions when costs and probabilities are both unknown.”, *Proceedings of the Seventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 204 - 213, 2001.

- [22] L. Breiman, “Bagging predictors.”, *Machine Learning*, vol. 24, no. 2, pp. 123 - 140, 1996.
- [23] R.E. Schapire, “The Strength of Weak Learnability (Extended Abstract).”, *Machine Learning*, vol. 5, pp. 197 - 227, 1990.
- [24] Y. Freund and R.E. Schapire, “A decision-theoretic generalization of on-line learning and an application to boosting.”, *Journal of Computer and System Sciences*, vol. 55, pp. 119 - 139, 1997.
- [25] M.R. Jung, Y.K. Park, J.W. Seon, K.Y. Kim, O. Cheong and S.Y. Ryu, “Definition and classification of complications of gastrectomy for gastric cancer based on the accordion severity grading system.”, *World Journal of Surgery*, vol. 36, no. 10, pp. 2400 - 2411, 2012.

저 자 소 개

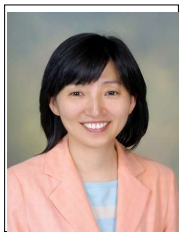


임수연(정회원)

2013년 전남대학교 간호학사 졸업.
2020년 전남대학교 간호학 석사 졸업.
2024년 전남대학교 간호학 박사 졸업.
2013년~2025년 화순전남대학교병원 간호사.
2024년~2025년 전남대학교 간호대학
시간강사.

2025년~현재 미시간대학교 간호대학
박사후연구원.

<주관심분야 : 중양간호, 인공지능, 빅데이터, 데이터
분석>



최자윤(정회원)

1991년 전남대학교 간호학사 졸업.
1994년 전남대학교 간호학 석사 졸업.
2000년 연세대학교 간호학 박사 졸업.
2005년~현재 전남대학교 간호대학 교
수.

<주관심분야 : 호흡기계 및 신경계 간호, 중환자 간호,
중양간호, 근거중심간호, 인공지능>