

머신러닝을 활용한 지역별 금융취약성지수의 구조적 기여 패턴 분석

임해리*¹, 송진아*², 강효연*³, 이종태**

Machine Learning-based Analysis of Structural Contribution Patterns in Regional Financial Vulnerability

Haeri Lim*¹, Jina Song*², HyoYeon Kang*³, and Jongtae Lee**

요 약

전국 단위 금융취약성지수(FVI)는 지역 간 경제·인구 구조의 이질성을 충분히 반영하지 못한다는 한계를 가진다. 이에 본 연구는 17개 광역자치단체를 대상으로 지역별 FVI를 재산출하고 머신러닝 기반 예측과 구조 분석을 수행하였다. 자산가격, 신용축적, 복원력 지표를 통합한 FVI를 종속변수로 설정하고 Random Forest, XGBoost, LightGBM 모델을 비교한 결과 XGBoost가 가장 우수한 예측 성능을 보였다. SHAP 분석과 K-means 군집분석을 통해 지역별 금융취약성 형성 구조의 유형 차이를 확인하였으며, PDP와 Shock 시나리오 분석은 거시 변수에 대한 비선형 반응과 지역별 민감도 차이를 보여준다. 특히 수도권은 고용 및 주거 안정 요인의 영향이 크고, 비수도권은 산업 및 인구 구조 요인의 영향이 두드러지는 것으로 나타났다. 이러한 결과는 획일적 정책 접근을 넘어 지역 특성을 반영한 금융안정 전략의 필요성을 시사한다.

Abstract

The national-level Financial Vulnerability Index (FVI) has limitations in reflecting regional heterogeneity in economic and demographic structures. This study reconstructs regional FVIs for 17 metropolitan and provincial governments in Korea and conducts machine-learning based prediction and structural analysis. Comparing Random Forest, XGBoost, and LightGBM models shows that XGBoost achieves the best predictive performance. SHAP analysis, K-means clustering, PDP, and shock scenario analysis reveal regional differences in the formation mechanisms and sensitivities of financial vulnerability. In particular, employment and housing stability factors are more influential in the Seoul metropolitan area, whereas industrial and demographic structures play a larger role in non-metropolitan regions. These findings suggest the need for financial stability strategies that account for regional characteristics.

Keywords

financial vulnerability index, regional financial risk, machine learning, SHAP, k-means clustering

* 서울여자대학교 데이터사이언스학과 학부과정
- ORCID¹: <https://orcid.org/0009-0003-4916-7581>
- ORCID²: <https://orcid.org/0009-0000-0903-5785>
- ORCID³: <https://orcid.org/0009-0004-4214-6873>
** 서울여자대학교 경영학과 부교수(교신저자)
- ORCID: <https://orcid.org/0000-0002-9685-7708>

• Received: Apr. 08, 2026 Revised: Jul. 26, 2026, Accepted: Jul. 29, 2026
• Corresponding Author: Jongtae Lee
Dept. of Business Administration, Seoul Women's University, Korea
Tel.: +82-2-970-5514, Email: light4u@swu.ac.kr

1. 서 론

금융취약성은 단일 금융지표의 일시적 변동이 아니라 신용 확대, 자산가격 상승, 거시경제 여건의 변화가 상호작용하며 축적되는 구조적 위험으로 이해된다. 금융위험을 측정하는 이론적 논의에서는 위험이 개별 변수의 선형적 변화로 즉각 드러나기보다는 누적적 과정 속에서 잠재적으로 증폭되며 일정 임계점에 도달할 때 표출될 수 있음을 강조해왔다[1]. 이는 금융취약성을 단순한 수준 지표가 아니라 시간에 따라 축적되는 동태적 구조로 해석해야 함을 시사한다.

이러한 문제의식은 국내에서도 제도화되었다. 금융취약성은 자산시장, 신용시장, 금융기관 건전성 등 금융시스템 전반에 내재된 위험요인이 누적된 상태를 의미하며, 이를 정량적으로 측정하기 위해 종합지수 형태의 금융취약성지수(FVI, Financial Vulnerability Index)가 편제되었다[2]. FVI는 특정 단일 변수의 변동을 반영하는 지표가 아니라 자산가격 과열, 신용 팽창, 레버리지 확대 등 금융불균형 요인을 통합적으로 포착하여 금융시스템의 잠재적 취약 정도를 나타내는 종합지표로 정의된다. 즉, 금융취약성지수는 금융불안이 발생하기 이전 단계에서 위험이 누적되는 구조적 상태를 측정하기 위한 지표라는 점에서 의미를 갖는다. 이는 금융취약성을 개별 변수의 합이 아닌 상호 연계된 구조로 이해해야 한다는 관점을 전제하고 있지만, 거시적 관점의 지표라는 한계를 가지고 있다.

유사하게, 금융시스템 전반의 스트레스를 측정하는 합성지표 접근에서도 시장 간 상관구조를 반영하여 통합지표를 구성함으로써 위험의 전이 가능성과 상호연계성을 고려하였다[3]. 이는 금융위험을 구조적 연결망 속에서 이해하려는 접근으로, 위험요인의 단순 나열이 아니라 기여 구조를 분석해야 한다는 점을 보여준다. 이러한 관점은 본 연구에서 변수 기여도를 기반으로 구조 유형을 도출하는 접근과 이론적 맥락을 공유한다.

한편 금융지표는 단순한 사후 진단 도구에 그치지 않고, 거시경제 변수와의 관계 속에서 사전적으로 예측 가능한 대상으로 다루어져 왔다. 전통적 계

량모형을 활용한 조기경보 접근에서는 금융불안 발생 가능성을 선제적으로 탐지하기 위한 확률모형을 제시하며 금융안정 모니터링이 예측 기반 분석으로 확장될 수 있음을 보였다[4]. 이는 금융취약성이 구조적 요인에 의해 설명되고 예측될 수 있다는 점을 뒷받침한다. 특히 금융취약성이 구조적·누적적 성격을 가진다는 점을 고려할 때, 지역 단위의 이질성을 반영한 분석은 금융안정 정책 측면에서도 중요한 의미를 가진다.

더 나아가 금융상황지수 및 금융스트레스지수의 구성과 활용을 다룬 연구에서는 지수 자체를 종속 변수로 설정하여 그 동학적 특성을 분석함으로써 지표가 단순한 현황 진단을 넘어 예측 대상이 될 수 있음을 실증하였다[5]. 즉, 금융지표는 정책 모니터링 수단일 뿐 아니라 거시경제 변수와의 상호작용 속에서 동태적으로 변화하는 분석 대상이라는 인식이 확립되고 있다.

그러나 기존 금융취약성 연구는 대체로 전국 단위 평균 지표를 중심으로 수행되어 지역 간 구조적 이질성을 충분히 반영하지 못한다는 한계를 가진다. 동일한 거시경제 충격이라 하더라도 지역의 산업구조, 인구구조, 주택시장 특성에 따라 금융위험의 형성 및 반응 양상은 상이하게 나타날 수 있다. 금융취약성은 다양한 거시경제 변수와 상호작용하며 비선형적으로 형성되는 특성을 가진다. 전통적 선형 회귀모형은 변수 간 복잡한 상호작용과 비선형 관계를 충분히 반영하는 데 한계가 존재한다.

이에 최근 금융위험 연구에서는 Random Forest, Gradient Boosting 계열 알고리즘 등 머신러닝 기반 접근이 활용되고 있으며, 복잡한 변수 구조와 비선형 패턴을 효과적으로 학습할 수 있다는 점에서 금융취약성 분석에도 적용 가능성이 확대되고 있다. 그러나 머신러닝 모형은 높은 예측 성능에도 불구하고 내부 의사결정 구조를 직접적으로 해석하기 어렵다는 한계를 가진다. 특히 금융취약성과 같이 정책적 활용 가능성이 중요한 분야에서는 단순한 예측 정확도뿐 아니라 어떠한 변수 구조가 금융취약성 형성에 기여하는지를 설명할 수 있는 해석 가능성이 요구된다.

이에 본 연구는 SHAP 기반 설명가능성 기법을

활용하여 변수별 기여 구조를 정량적으로 분석하고, 지역별 금융취약성 형성 패턴의 차이를 구조적으로 해석하고자 한다. 보다 구체적으로, 본 연구는 기존 금융취약성지수(FVI)의 개념적 틀을 지역 단위로 확장하고, 머신러닝 기반 구조 학습과 설명가능성 기법을 결합하여 지역별 금융취약성이 어떠한 거시 경제 및 인구구조 변수와 연관되어 형성되는지를 분석하여 동일한 거시경제 변수라도 지역에 따라 상이한 기여 구조가 나타나는 가를 확인한다. 또한 SHAP 기반 지역 프로파일과 군집 분석, Shock 시나리오 분석을 결합함으로써 지역 금융취약성의 구조 유형과 외생 충격에 대한 반응 차이를 실증적으로 비교하고자 한다.

본 연구의 차별성은 다음과 같다. 첫째, 전국 단위 금융취약성지수를 지역 단위로 확장함으로써 지역 간 구조적 이질성을 정량적으로 반영하였다. 둘째, 단순한 예측 성능 비교를 넘어 SHAP 기반 설명가능성 분석을 활용하여 금융취약성의 형성 구조를 해석하였다. 셋째, SHAP 기반 지역 프로파일과 군집 분석, Shock 시나리오 분석을 결합하여 구조 유형과 충격 반응 간의 연계성을 실증적으로 제시하였다.

본 논문의 구성은 다음과 같다. 2장에서는 금융취약성 및 머신러닝 기반 예측 연구에 대한 선행연구를 검토한다. 3장에서는 데이터 구성 및 연구 방법론을 설명하고, 4장에서는 머신러닝 기반 분석 결과와 구조 유형 분석 결과를 제시한다. 마지막으로 5장에서는 연구 결과를 요약하고 정책적 시사점 및 한계를 논의한다.

II. 관련 선행 연구

2.1 금융취약성 및 위험 예측 연구

금융취약성(Financial vulnerability)과 시스템 리스크는 글로벌 금융위기 이후 거시건전성 정책의 핵심 연구 주제로 부각되었다. 초기 연구들은 금융시장 불안정성을 종합적으로 측정하기 위하여 자산가격 변동성, 신용스프레드, 환율 변동성 등을 결합한 금융스트레스지수(FSI, Financial Stress Index)를 제안

하였다[6]. 이러한 접근은 금융불안의 축적 과정을 단일 변수로 설명하기보다 복수의 거시·금융 지표를 종합적으로 반영해야 함을 제시하였다. 그러나 해당 연구들은 주로 국가 전체 수준의 금융불안 측정에 초점을 두고 있어 지역별 구조적 차이를 반영하는 데에는 한계가 존재한다.

이와 더불어 신용갭(Credit-to-GDP gap), 자산가격 상승률 등 거시금융 변수를 활용하여 위기 발생 가능성을 확률적으로 추정하는 조기경보모형(Early warning model)이 제안되었다[7][8]. 해당 연구들은 신용팽창과 자산시장 과열이 금융위기 발생 확률과 통계적으로 유의한 관계를 가진다는 점을 실증적으로 보였으며 패널 데이터 기반 예측 모형을 활용하여 국가 간 위험 축적 패턴을 비교하였다. 다만 위기 발생 가능성 자체의 예측에 초점을 두고 있어 금융취약성이 어떠한 구조적 경로를 통해 형성되는지에 대한 해석은 상대적으로 제한적이었다.

국내에서는 가구 단위 미시자료를 활용하여 지역별 가계부실 위험을 추정하고, 소득·부채·자산구조가 부실확률에 미치는 영향을 분석한 연구가 수행되었다[9]. 해당 연구는 지역 간 위험 격차를 실증적으로 제시하였다는 점에서 의의가 있으나 개별 가구의 부실위험 예측에 초점을 두고 있어 지역 거시구조 차원에서의 종합적 금융취약성 형성 메커니즘을 분석하는 데에는 한계가 존재한다.

최근에는 전통적 회귀모형을 넘어 머신러닝 기법을 활용하여 금융위기 또는 신용위험을 예측하려는 연구가 증가하고 있다. Random Forest, Gradient Boosting 등 비선형 모형은 복합적 거시변수 간 상호작용을 효과적으로 학습할 수 있다는 점에서 금융 리스크 예측에 활용되고 있으며[10][11] 일부 연구에서는 머신러닝 모형이 전통적 로짓 모형 대비 높은 예측 성능을 보인다고 보고하였다[12]. 그러나 이러한 연구들은 높은 예측 정확도를 달성하는 데 초점을 두는 경우가 많으며, 머신러닝 모형 내부의 변수 기여 구조나 지역별 구조적 이질성을 해석하려는 접근은 상대적으로 제한적이다.

이러한 선행 연구들은 주로 전국 단위 평균 지표를 활용하거나 금융위기 발생 확률 예측에 초점을 두고 있어 지역별 구조적 이질성을 충분히 반영하

지 못한다는 한계가 존재한다. 또한 설명가능성 기반 구조 분석보다는 예측 정확도 개선 자체에 집중하는 경향이 강하다. 이에 본 연구는 지역 단위 금융취약성을 대상으로 머신러닝 예측과 설명가능성 분석을 결합하여 구조적 차이를 분석하고자 한다.

2.2 트리 기반 앙상블 예측 모형

2.2.1 Random Forest

Random Forest는 Breiman[10]에 의해 제안된 트리 기반 앙상블 학습 방법으로, 다수의 결정트리를 결합하여 일반화 오차를 감소시키는 것을 목적으로 한다. 해당 모형은 Bootstrap Aggregating(Bagging)을 기반으로 하며 서로 다른 부트스트랩 표본에서 학습된 다수의 트리를 평균화함으로써 예측의 분산을 줄이는 구조를 가진다. 주어진 입력 x 에 대해, B 개의 결정트리 $h_b(x)$ 가 생성될 때 회귀 문제에서의 최종 예측값은 식 (1)과 같이 정의된다.

$$\widehat{f(x)} = \frac{1}{B} \sum_{b=1}^B h_b(x) \quad (1)$$

이와 같이 개별 트리의 예측을 평균화함으로써 모형의 분산을 감소시키는 것이 Random Forest의 기본 원리이다. 그러나 단순한 Bagging과 달리 Random Forest는 각 노드 분할 시 전체 변수 집합이 아닌 무작위로 선택된 일부 변수만을 고려하여 최적 분할을 수행한다. 이러한 무작위 특성 선택(Random feature selection)은 트리 간 상관관계를 낮추는 역할을 하며 결과적으로 앙상블 모형의 일반화 오차를 추가적으로 감소시킨다.

Breiman[10]은 Random Forest의 성능이 개별 트리의 예측 강도와 트리 간 상관관계에 의해 결정된다는 점을 이론적으로 제시하였다. 즉, 강한 예측력을 가진 트리들을 유지하면서도 상호 상관성을 낮추는 구조가 전체 모형의 안정성과 예측력을 동시에 확보하는 핵심 메커니즘이다. 또한 트리 개수가 충분히 증가할 경우 일반화 오차가 수렴하며 과적합이 발생하지 않는다는 점도 제시되었다.

이와 같은 구조적 특성에 기반하여, Random Forest는 복수의 설명변수 간 비선형 관계와 상호작용을 안정적으로 학습할 수 있는 앙상블 모형으로 간주되며, 본 연구에서는 금융취약성의 복합적 결정 구조를 포착하기 위한 비교 기준 모형으로 채택하였다.

2.2.2 Gradient Boosting

Gradient Boosting은 Friedman[13]에 의해 제안된 가산적 함수 추정(Additive function approximation) 기반 앙상블 학습 방법으로, 약한 학습기(Weak learner)를 순차적으로 결합하여 손실함수를 최소화하는 방향으로 모델을 확장한다. Gradient Boosting의 기본 구조는 식 (2)와 같은 가산 모형으로 표현된다.

$$F_{M(x)} = \sum_{m=1}^M \gamma_m h_m(x) \quad (2)$$

여기서 $h_m(x)$ 는 각 단계에서 학습되는 결정트리이며, γ_m 은 해당 트리의 기여도를 나타낸다. 각 단계에서 현재 모델은 손실함수 $L(y, F(x))$ 에 대해 음의 그레디언트를 계산하고, 이를 근사하는 트리를 학습함으로써 모델을 점진적으로 계산한다. 즉, 이전 단계의 예측 오차를 보정하는 방식으로 편향(Bias)을 감소시키는 구조를 가진다. XGBoost는 Gradient Boosting 프레임워크를 확장하여 정규화 항을 포함한 목적함수를 도입함으로써 모델 복잡도를 명시적으로 통제하는 구조를 가진다[11]. XGBoost의 목적함수는 식 (3)과 같이 정의된다.

$$L^{(t)} = \sum_{i=1}^n l(y_i, \hat{y}_i^{(t-1)} + f_{t(x_i)}) + \Omega(f_t) \quad (3)$$

여기서 $l(\cdot)$ 은 손실함수, f_t 는 t 번째 트리, $\Omega(f_t)$ 는 트리 복잡도에 대한 정규화 항이다. 정규화 항은 식 (4)와 같이 표현된다.

$$\Omega(f) = \gamma T + \frac{1}{2} \lambda \sum_{j=1}^T w_j^2 \quad (4)$$

여기서 T 는 리프 노드 수, ω 는 각 리프의 가중치이다. 이 구조는 트리의 복잡성을 직접적으로 패널티화하여 과적합을 억제하는 역할을 한다. 또한 XGBoost는 목적함수를 2차 테일러 전개(Second-order Taylor expansion)로 근사하여 그래디언트와 헤시안 정보를 동시에 활용함으로써 분할 기준을 정교하게 계산한다. 이러한 2차 정보 기반 최적화는 학습 안정성과 수렴 속도를 동시에 향상시키는 구조적 특징을 가진다.

LightGBM은 GBDT(Gradient Boosting Decision Tree)를 대규모 데이터 환경에서 효율적으로 구현하기 위해 설계된 알고리즘으로, 히스토그램 기반 분할(Histogram-based split)과 leaf-wise 트리 성장 전략을 핵심으로 한다[14]. LightGBM 식 (5)의 손실 최소화 문제를 기반으로 학습을 수행한다.

$$\min \sum_{i=1}^n L(y_i, F(x_i)) \quad (5)$$

그러나 기존 level-wise 방식과 달리, 손실 감소가 가장 큰 리프를 우선적으로 분할하는 leaf-wise 성장 전략을 채택한다. Leaf-wise 방식은 동일한 트리 깊이에서도 더 큰 손실 감소를 달성할 수 있으나, 과도한 분할을 방지하기 위해 깊이 제한(Max depth) 등의 제약 조건이 병행된다. 또한 히스토그램 기반 분할은 연속형 변수를 구간화하여 계산 복잡도를 감소시키는 구조를 가진다.

Gradient Boosting 계열 모형은 이와 같이 가산적 함수 추정 구조에 기반하여 손실함수를 직접 최소화하는 방향으로 모델을 확장하며, XGBoost와 LightGBM은 각각 정규화 기반 복잡도 통제와 효율적 분할 전략을 통해 예측 안정성과 계산 효율성을 강화한 확장 형태로 볼 수 있다. 이러한 구조적 특성에 기반하여, 본 연구에서는 금융취약성의 복합적 비선형 관계를 정교하게 학습하기 위한 핵심 예측 모형으로 해당 계열 알고리즘을 채택하였다.

2.3 모델 해석 기법 연구

SHAP(SHapley Additive exPlanations)는 협력적 게임이론(Cooperative game theory)에 기반하여 예측모

형의 개별 변수 기여도를 정량적으로 분해하는 방법이다[15][16]. SHAP은 Shapley value 개념을 머신러닝 예측모형에 적용하여 특정 예측값이 각 설명 변수에 의해 어떻게 형성되었는지를 일관된(Additive) 형태로 표현한다. 주어진 예측모형 $f(x)$ 에 대해 SHAP은 식 (6)과 같은 가산적 설명 모형(Additive feature attribution model)을 가정한다. 식 (6)은 개별 예측값이 기준값(ϕ_0)과 각 변수의 SHAP 기여도(ϕ_i)의 합으로 분해될 수 있음을 나타낸다.

$$f(x) = \phi_0 + \sum_{i=1}^M \phi_i \quad (6)$$

여기서 ϕ_0 는 전체 데이터에 대한 평균 예측값, ϕ_i 는 각 변수 i 의 기여도(Shapley value)를 의미한다. 즉, SHAP은 개별 예측값을 기준값(Baseline)과 변수별 기여도의 합으로 분해한다. Shapley value는 모든 변수 조합에 대한 기여도의 평균적 한계 기여(Marginal contribution)를 계산하여 정의된다. 변수 i 의 Shapley value는 식 (7)과 같이 표현된다. 식 (7)은 특정 변수가 예측 결과 형성에 기여하는 평균 한계효과를 모든 변수 조합에 대해 계산하는 과정을 의미한다.

$$\phi_i = \sum_{S \subseteq M} \frac{|S|! (M - |S| - 1)!}{M!} [f(S \cup i) - f(S)] \quad (7)$$

여기서 S 는 변수 부분집합, M 은 전체 변수 개수이다. 해당 식은 변수 i 가 모든 가능한 변수 조합에 추가될 때의 평균적 기여도를 계산하는 구조를 가진다. Lundberg & Lee는 이러한 Shapley value가 국소 정확성(Local accuracy), 결측 무기여성(Missingness), 일관성(Consistency)의 세 가지 공리적 특성을 만족함을 제시하였다. 이로 인해 SHAP은 다양한 설명 기법 중 유일하게 이론적으로 일관된 변수 기여도 분해를 제공하는 방법으로 평가된다.

특히 트리 기반 앙상블 모형에 대해서는 TreeSHAP 알고리즘이 제안되어 모든 변수 조합을 완전 탐색하지 않고도 정확한 Shapley 값을 효율적

으로 계산할 수 있음이 제시되었다[15]. 이는 Random Forest 및 Gradient Boosting 계열 모형에 대해 계산 가능성과 해석 가능성을 동시에 확보하는 이론적 근거가 된다.

이와 같은 구조적 특성에 기반하여 SHAP은 복잡한 비선형 예측모형의 내부 구조를 일관된 기여도 형태로 분해할 수 있는 방법론으로 간주되며, 본 연구에서는 금융취약성 예측 결과의 변수별 영향 구조를 정량적으로 해석하기 위한 핵심 도구로 채택하였다.

최근에는 설명 가능한 인공지능(XAI)을 활용하여 머신러닝 모형의 예측 결과를 해석하려는 연구가 금융위험 분야에서도 활발히 이루어지고 있다. 이러한 연구들은 단순한 예측 성능 평가를 넘어 개별 변수들이 예측값 형성에 어떠한 기여 패턴을 보이는지를 분석함으로써 모형의 해석 가능성과 신뢰성을 향상시키고자 한다. 특히 SHAP 기반 접근은 신용위험 및 부도예측 분야에서 주요 예측 변수의 상대적 기여도를 정량적으로 분석하고, 금융기관의 의사결정 과정에 대한 설명력을 제공하는 방법론으로 활용되고 있다. 또한 최근 연구에서는 설명가능성 기법이 금융위험 예측 결과의 투명성을 높이고 규제 및 정책 의사결정을 지원하는 도구로 활용될 수 있음이 보고되고 있다[17][18].

2.4 군집 분석 기반 구조 유형화 연구

K-means 군집분석은 MacQueen[19]에 의해 제안된 비지도 학습 기법으로, 주어진 데이터 집합을 사전에 정의된 K 개의 군집으로 분할하여 군집 내 분산(Within-cluster variance)을 최소화하는 것을 목적으로 한다. 해당 방법은 데이터의 구조적 유사성을 거리 기반으로 측정하며 각 관측치를 가장 가까운 중심점(Centroid)에 할당하는 반복적 최적화 과정을 수행한다. 주어진 관측치 집합 $\{x_1, x_2, \dots, x_n\}$ 에 대해, K-means는 식 (8)의 목적함수를 최소화하는 방향으로 군집을 형성한다. 식 (8)은 각 관측치와 소속 군집 중심 간 거리 제곱합을 최소화하는 과정을 나타낸다.

$$\min_{C_1, \dots, C_K} \sum_{k=1}^K \sum_{x_i \in C_k} \|x_i - \mu_k\|^2 \quad (8)$$

여기서 C_k 는 k 번째 군집, μ_k 는 해당 군집의 중심점을 의미한다. K-means는 각 관측치를 최근접 중심점에 할당하고, 군집 내 관측치들의 평균으로 중심점을 갱신하는 과정을 반복함으로써 군집 내 제곱거리합(Within-cluster sum of squares)을 최소화한다[20]. 이 반복적 최적화 과정은 목적함수가 수렴할 때까지 수행되며, 결과적으로 군집 내 응집도를 극대화하는 방향으로 군집 구조를 형성한다.

군집 중심이 평균 벡터로 정의된다는 점에서, K-means는 연속형 변수로 구성된 다변량 벡터 공간에서 거리 기반 유사성을 측정하는 데 적합한 방법론으로 간주된다. 따라서 각 관측치를 구조적 특성을 반영하는 벡터로 표현할 수 있는 경우, 해당 벡터들 간의 유사성에 기반한 유형화가 가능하다.

이와 같은 이론적 특성에 기반하여, K-means는 다변량 변수로 구성된 구조 벡터를 유형화하는 데 적합한 방법론으로 간주되며, 본 연구에서는 지역별 영향 구조의 유사성을 기반으로 유형을 분류하기 위한 군집 알고리즘으로 채택하였다.

III. 연구 설계

3.1 연구 방법론 설계

기존의 금융취약성지수(FVI) 관련 연구는 지수 산출 방식의 개선이나 예측 정확도 향상에 초점을 두어왔다. 특히 머신러닝을 활용한 연구의 경우 예측 성능 비교에 집중하는 경향이 있으며, 개별 거시경제 변수가 금융취약성에 어떠한 구조적 관계를 형성하는지에 대한 체계적 분석은 상대적으로 부족한 실정이다.

이에 본 연구는 예측 정확도 중심 접근을 보완하고, 지역별 금융취약성의 결정 구조를 설명하기 위한 분석 프레임워크를 제안한다. 제안하는 프레임워크는 머신러닝 기반 구조 모형 학습과 설명가능성 분석을 결합하여, 거시경제 변수와 금융취약성 간의 관계를 정량적으로 도출하고 지역 간 구조적 차이를 비교하는 것을 목표로 한다.

본 연구의 분석 절차는 다음과 같이 구성된다. 첫째, 지역별 패널 데이터를 구축하고 금융취약성의

시차 구조를 포함한 설명 모형을 설계한다. 둘째, 트리 기반 머신러닝 모형을 학습하여 비선형 관계 및 변수 간 상호작용을 반영한다. 셋째, SHAP을 활용하여 변수별 기여도를 산출하고, 이를 통해 글로벌 수준과 지역 수준의 영향 구조를 분석한다. 넷째, 부분의존도분석(PDP, Partial Dependence Plot)을 통해 주요 변수와 금융취약성 간의 평균적 관계를 시각적으로 검토한다. 다섯째, SHAP 기반 지역 프로파일을 활용하여 군집화를 수행함으로써 금융취약성 결정 구조의 유사성에 따라 지역을 유형화한다. 마지막으로, 주요 거시경제 변수에 대한 충격 시나리오를 설정하고 군집별 반응 차이를 비교한다.

3.2 데이터 구성

본 연구는 2004년부터 2024년까지 20개 연도에 걸친 17개 광역자치단체 패널 데이터를 활용하였다. 분석 대상 지역은 서울특별시, 부산광역시, 대구광역시, 인천광역시, 광주광역시, 대전광역시, 울산광역시, 세종특별자치시, 경기도, 강원특별자치도, 충청북도, 충청남도, 전북특별자치도, 전라남도, 경상북도, 경상남도, 제주특별자치도로 구성된다.

종속변수는 지역별 금융취약성지수(FVI)이다. FVI는 자산가격(Asset), 신용축적(Credit), 복원력(Resilience)의 세 부분을 통합하여 금융시스템의 구조적 취약 수준을 측정하는 지표이다.

자산가격 부문은 주택가격지수, 주택매매가격, 주택가격 대비 가구소득비율(PIR, Price-to-Income Ratio) 등을 활용하여 부동산 및 자산시장 과열 정도를 반영한다. 신용축적 부문은 가계신용증가율, 기업대출증가율, 예금은행대출금, 단기외화부채비율, 원화대출금비율 등을 활용하여 지역 내 신용팽창과 부채 누적 수준을 측정한다.

복원력 부문은 지역 단위 금융기관 건전성 자료의 확보가 제한된다는 점을 고려하여, 지역 경제·인구 구조가 금융충격에 대응하고 회복할 수 있는 능력을 나타내는 대리변수(Proxy variables)로 구성하였다. 구체적으로 실업률, 고령화율, 1인당 지역내총생산(GRDP, Gross Regional Domestic Product), 평균 가구원수, 자가주택 비율 및 임차가구 비율을 활용하

여 지역의 경제적 안정성과 구조적 회복력을 측정하였다.

즉, 본 연구의 복원력 부문은 한국은행 FVI의 복원력 개념을 지역 단위로 확장하여 지역경제 복원력(Regional resilience)의 관점에서 재구성한 것이다. 본 연구는 이러한 세 부분을 통합함으로써 지역 금융취약성을 단일 변수 수준이 아닌 구조적 위험의 관점에서 측정하고자 하였다.

각 부문을 구성하는 변수들은 취약성 방향을 통일한 후 식 (9)의 Min - Max Scaling을 적용하여 0~1 구간으로 정규화하였다. 이는 변수 간 단위 차이와 규모 효과를 제거하고 서로 다른 단위 체계를 비교 가능한 형태로 변환하기 위한 과정이다. 정규화 식은 다음과 같다.

$$Z_{i,t,k} = \frac{X_{i,t,k} - \min(X_{i,t,k})}{\max(X_{i,t,k}) - \min(X_{i,t,k})} \quad (9)$$

여기서 i 는 지역, t 는 연도, k 는 변수를 의미한다. 식 (9)의 $X_{i,t,k}$ 는 특정 연도 t 에서 변수 k 에 대해 관측된 전체 지역의 값 집합을 의미하며 $\min(X_{i,t,k})$ 와 $\max(X_{i,t,k})$ 는 각각 동일 연도와 동일 변수 기준에서 전체 지역 중 최솟값과 최댓값을 의미한다. 따라서 식 (9)는 지역 i 의 변수값 $X_{i,t,k}$ 를 해당 연도·변수의 지역 간 분포 내 상대적 위치로 변환하는 정규화 과정이다. 이를 통해 서로 다른 단위를 갖는 변수들을 동일 척도 상의 상대적 취약성 지표로 변환하였다. 정규화된 변수들을 부문별로 평균하여 각 부문 지수를 산출한 후, 식 (10)과 같이 세 부문의 단순 평균을 통해 연도별 지역 FVI 원지수($FVI_{i,t}^{raw}$)를 계산하였다. 이는 개별 변수의 단기 변동보다 자산가격, 신용축적, 복원력 차원의 구조적 위험 수준을 종합적으로 반영하기 위한 목적이다.

$$FVI_{i,t}^{raw} = \frac{1}{3}(P_{i,t}^{(Asset)} + P_{i,t}^{(Credit)} + P_{i,t}^{(Resilience)}) \quad (10)$$

이후 전국 단위 FVI와의 수준 차이를 보정하기 위하여 연도별 보정계수 S_t 를 산출하고, 식 (11)을 이용하여 보정된 최종 지역 금융취약성지수(FVI)를 구성하였다.

$$FVI_{i,t}^{cal} = FVI_{i,t}^{raw} \times S_i \quad (11)$$

보정계수는 동일 연도의 전국 평균 지역 FVI_raw와 한국은행 FVI 간의 상대적 비율을 기준으로 산출하였으며, 이를 통해 지역별 상대적 구조를 유지하면서 전국 단위 금융취약성 수준과의 정합성을 확보하고자 하였다. 이를 통해 전국 기준과 정렬된 지역 간 상대적 금융취약 수준을 확보하였다. 설명변수는 지역의 거시경제 및 인구 구조 특성을 반영하는 동적 변수로 구성하였다. 구체적으로 고령화율, 실업률, 1인당 지역내총생산(GRDP per capita), 평균 가구원 수, 자가 점유 비율, 임차 비율을 포함하였다.

고령화율은 전체 인구 대비 65세 이상 인구 비율로 산출하였으며, 실업률은 지역별 연평균 실업률을 활용하였다. 1인당 GRDP는 지역내총생산을 해당 지역 인구로 나누어 계산하였다. 평균 가구원 수 및 주택 점유 형태는 인구주택총조사 자료를 기반으로 산출하였다. 모든 통계자료는 국가통계포털(KOSIS)을 통해 수집하였다. 본 연구는 지역 경제·인구 구조의 시간적 변동을 반영하기 위하여 각 변수의 연도별 값을 그대로 활용하였다.

금융취약성의 자기상관 구조를 반영하기 위하여 FVI의 1기, 2기, 3기 시차 변수($FVI_{i,t-1}$, $FVI_{i,t-2}$, $FVI_{i,t-3}$)를 추가하였다. 이는 단기 관성 효과와 누적적 위험 구조를 모형에 포함하기 위함이다. 일부 연도에 존재하는 결측치는 시계열의 연속성을 유지하기 위해 선형 보간법을 적용하였다. 최종적으로 분석 데이터는 지역-연도 단위의 균형 패널 형태로 구성되었으며 시차 변수 적용에 따라 초기 일부 연도는 분석에서 제외되었다.

3.3 구조 모형 설계

본 연구는 지역별 금융취약성지수(FVI)의 구조적 결정 요인을 분석하기 위하여 트리 기반 머신러닝 모형을 활용하였다. 본 연구의 목적은 단순한 예측 정확도의 향상이 아니라, 거시경제 환경 변화에 따라 금융취약성이 어떠한 구조로 반응하는지를 규명하는 데 있다. 이에 따라 변수 간 비선형 관계와 상호작용 효과를 동시에 포착할 수 있는 비모수적 모

형을 채택하였다.

금융취약성은 시계열적 자기상관 특성을 가지는 지표로, 특정 시점의 취약성 수준은 이전 시점의 상태에 의해 부분적으로 결정된다. 이러한 관성 효과를 반영하기 위하여 본 연구는 FVI의 1기, 2기, 3기 시차값을 설명변수에 포함하였다. 이를 통해 현재 시점의 금융취약성은 과거 취약성의 누적적 영향과 동시점 거시경제 요인의 결합 결과로 형성된다는 구조적 가정을 반영하였다. 즉, 본 연구는 식 (12)와 같은 함수 형태를 전제로 한다.

$$FVI_{i,t} = f(FVI_{i,t-1}, FVI_{i,t-2}, FVI_{i,t-3}, X_{i,t}) \quad (12)$$

여기서 $X_{i,t}$ 는 지역 i 의 t 시점 거시경제 및 인구구조 변수 벡터를 의미하며, 고령화율, 실업률, 1인당 GRDP, 평균 가구원수, 주택 점유형태 비율 등을 포함한다. 또한 함수 $f(\cdot)$ 는 변수 간 관계에 대해 사전적 함수 형태를 가정하지 않는 비모수적 접근을 의미한다.

한편, 본 연구에서는 시간 인덱스(Year) 변수를 명시적으로 포함하지 않았다. 이는 장기 추세를 직접 학습하는 시간 변수가 구조적 요인의 상대적 기여도를 왜곡할 가능성을 배제하기 위함이다. 본 연구의 관심은 시간 경과 자체가 아니라 동일 시점의 거시경제 조건 하에서 금융취약성이 어떻게 결정되는가에 있으므로, 시간 효과는 시차 변수에 내재된 동학 구조를 통해 간접적으로 반영되도록 설계하였다.

모형으로는 RandomForest, LightGBM, XGBoost를 비교 대상으로 설정하였다. 이들 알고리즘은 모두 결정트리를 기반으로 하는 앙상블 학습 기법으로, 변수 간 비선형성과 고차 상호작용을 명시적 모형 가정 없이 추정할 수 있다는 장점을 가진다. 알고리즘 간 성능 비교를 통해 최적 모형을 선정한 후, 해당 모형을 기반으로 구조 해석 분석을 수행하였다.

학습 및 검증 과정에서는 랜덤 분할 방식을 적용하되, 지역별 표본 비율이 유지되도록 분할하였다. 이는 특정 지역에 표본이 편중되는 현상을 방지하고, 모형이 지역 특성에 과도하게 적합되는 것을 방지하기 위함이다.

3.4 해석 및 유형화 설계

본 연구는 머신러닝 모델을 통해 추정된 결과를 해석하기 위하여 설명가능성 기법을 적용하였다. 트리 기반 모형은 높은 예측 성능을 보이는 반면, 내부 의사결정 구조가 직접적으로 해석되기 어렵다는 한계를 가진다. 이에 따라 본 연구는 변수 기여도를 정량적으로 산출할 수 있는 SHAP(Shapley Additive Explanations) 기법을 활용하였다. SHAP은 게임이론의 Shapley value 개념에 기반하여 각 설명변수가 예측값에 기여한 한계 공헌도를 계산하는 방법이다. 본 연구에서는 이러한 SHAP 값을 활용하여 변수별 전반적 영향 구조와 지역별 기여도 차이를 분석하였다.

또한 주요 거시경제 변수와 금융취약성 간의 관계를 직관적으로 파악하기 위하여 부분의존도분석(PDP)을 수행하였다. PDP는 특정 변수를 일정 범위에서 변화시킬 때 다른 변수의 평균 효과를 통제된 상태에서 예측값이 어떻게 변하는지를 나타내는 방법으로, 변수와 종속변수 간의 평균적 부분 효과를 시각화한다[13].

이후, 지역 간 구조적 유사성을 파악하기 위하여 SHAP 기반 지역별 평균 기여도 값을 활용한 K-means 군집화를 수행하였다. 이는 행정구역 기준이 아닌 금융취약성 결정 구조의 유사성에 기반하여 지역을 유형화하기 위한 절차이다. 군집 개수는 실루엣 지수(Silhouette score)를 기준으로 결정하였다.

마지막으로, 거시경제 변수의 외생적 변동이 금융취약성에 미치는 영향을 정량적으로 평가하기 위하여 충격 시나리오 분석을 수행하였다. 각 변수의 분포를 고려하여 표준편차(1σ) 수준의 변화를 충격 크기로 설정하고, 그에 따른 예측 FVI의 평균 변화를 산출하였다. 이는 변수 간 단위 및 변동 폭의 차이를 보정한 상태에서 상대적 영향력을 비교하기 위함이다. 또한 이러한 변동 효과를 군집별로 산출하여 구조 유형에 따른 반응 차이를 검토하였다.

IV. 연구 결과

4.1 예측 모형 성능 평가

본 연구에서는 금융취약성지수(FVI)의 구조적 결정요인을 분석하기 위하여 RandomForest, LightGBM, XGBoost 모형을 적용하고 그 성능을 비교하였다. 비교는 시간 추세 변수를 제외한 모형을 기준으로 수행되었으며, 평가지표로는 결정계수(R^2), 평균절대오차(MAE), 평균제곱근오차(RMSE)를 활용하였다.

분석 결과, XGBoost 모형이 가장 우수한 설명력을 보였다. 구체적으로 XGBoost의 R^2 는 0.843으로 나타났으며, 이는 FVI 변동의 약 84.3%가 전기값 및 거시경제 변수에 의해 설명될 수 있음을 의미한다. 평균절대오차(MAE)는 4.545, 평균제곱근오차(RMSE)는 8.057로 나타나, 예측 오차 역시 안정적인 수준임을 확인하였다. LightGBM과 RandomForest 또한 유사한 수준의 성능을 보였으나, 전반적인 설명력과 오차 지표 측면에서 XGBoost가 가장 우수한 결과를 나타냈다. 이에 따라 본 연구에서는 이후의 구조 해석 분석을 XGBoost 모형을 중심으로 수행하였다(표 1).

표 1. 머신러닝 모델 예측 성능 비교

Table 1. Comparison of predictive performance across machine learning models

model	R^2	MAE	RMSE
RandomForest	0.80432	4.76887	9.00388
LightBM	0.83977	5.50600	8.14743
XGBoost	0.84329	4.54544	8.05746

Random Split 기반 검증의 한계를 보완하기 위하여 연도 기준 Expanding Window 방식의 Panel Cross Validation을 수행하였다. 이는 과거 연도 자료를 학습 데이터로 사용하고 이후 연도 자료를 검증 데이터로 사용하는 방식으로, 시계열적 순서를 고려한 보다 엄격한 검증 절차에 해당한다. 분석 결과 모든 모형에서 Random Split 대비 예측 성능이 낮게 나타났으며, 이는 미래 시점 예측의 어려움과 코로나19와 같은 구조적 충격의 영향을 반영하는 것으로 해석된다. 그럼에도 XGBoost 모형은 Panel Cross Validation에서도 $R^2=0.5747$, $RMSE=9.1536$ 을 기록하여 비교 모형 중 상대적으로 우수한 성능을 나타냈다. 따라서 본 연구는 예측 성능과 일반화 가능성을

종합적으로 고려하여 이후 분석을 XGBoost 모형을 중심으로 수행하였다.

이와 함께, 과적합 여부를 검토하기 위하여 학습 데이터와 검증 데이터의 성능을 비교하였다. 비교 결과 일부 모형에서 학습 성능이 검증 성능보다 높게 나타나 과적합 가능성이 확인되었다. 이에 트리 깊이(max_depth) 제한, 트리 수(n_estimators) 조정, 정규화 항(reg_lambda) 적용 등 모델 복잡도를 제어하는 설정을 적용하였다. 그럼에도 17개 지역 × 20개 연도의 제한된 패널 자료를 활용한 데 따른 한계로 학습 성능과 검증 성능 간 차이는 일부 잔존하는 것으로 나타났다(표 2).

표 2. Panel Cross Validation 기반 강건성 검증 결과
Table 2. Robustness check results based on panel cross validation

model	R ²	MAE	RMSE
RandomForest	0.5330	7.7301	9.6603
LightGBM	0.5373	7.2740	9.4232
XBoost	0.5747	7.0752	9.1536

4.2 SHAP 기반 금융취약성 결정 요인의 구조 분석

본 연구에서는 XGBoost 모형을 대상으로 SHAP(SHapley Additive exPlanations) 분석을 수행하여 각 설명변수가 금융취약성지수(FVI)에 미치는 기여도를 정량적으로 평가하였다. SHAP 값은 개별 관측치에 대해 각 변수가 예측값에 기여한 정도를 계산하며, 양(+)의 값은 FVI를 증가시키는 방향의 기여를, 음(-)의 값은 감소시키는 방향의 기여를 의미한다.

Global SHAP 중요도 분석 결과, 가장 예측 기여도가 높게 나타난 변수는 FVI의 시차 변수(FVI_lag1, FVI_lag2, FVI_lag3)로 나타났다. 특히 FVI_lag1의 평균 절대 SHAP 값은 13.04로 다른 변수들에 비해 압도적으로 높게 나타나, 금융취약성의 지속성(persistence)이 강하게 반영되는 경향을 보였다. 이는 지역 금융취약성의 현재 수준이 과거 취약성 수준과 높은 연관 구조를 가지는 특성을 시사한다(그림 1, 그림 2).

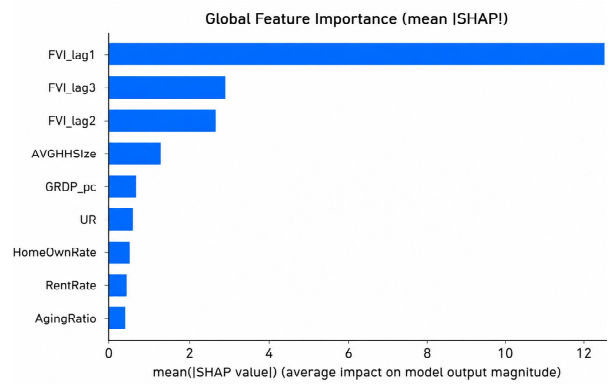


그림 1. 글로벌 SHAP 중요도

Fig. 1. Global SHAP feature importance of the prediction model

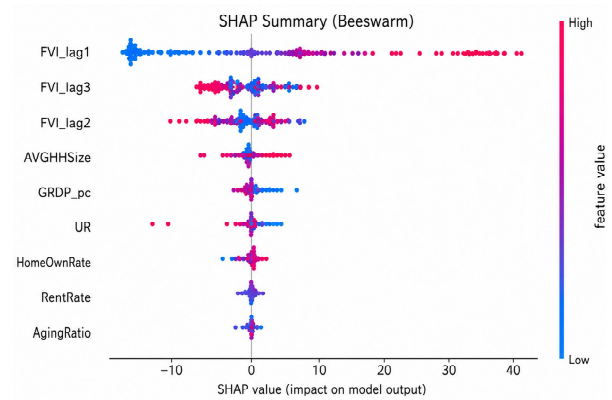


그림 2. 예측 모형의 변수 중요도 및 영향 방향

Fig. 2. Feature importance and effect direction of the prediction model

거시경제 변수 중에서는 평균가구원수, 1인당 GRDP, 자가주택비율, 실업률, 고령화율, 임차비율 순으로 예측 기여도가 나타났다. 평균가구원수와 1인당 GRDP는 상대적으로 높은 SHAP 값을 보여 지역의 인구구조 및 소득 수준이 금융취약성에 중요한 구조적 요인임을 시사한다.

한편, 지역별 SHAP 분석 결과 동일한 변수라도 지역에 따라 기여 방향이 상이하게 나타났다. 예를 들어 일부 지역에서는 1인당 GRDP가 FVI를 증가시키는 방향으로 작용한 반면, 다른 지역에서는 감소 방향으로 작용하였다.

이는 동일한 경제 변수라도 지역의 산업구조, 인구구성, 주택시장 특성 등에 따라 금융취약성과의 연관 구조가 다르게 나타날 가능성을 시사한다. 또한 일부 지역에서는 고령화율이 음(-)의 SHAP 값을 보이며 FVI 예측값을 완화하는 방향의 기여를

나타낸 반면, 다른 지역에서는 양(+)의 방향으로 나타났다. 이는 고령화율과 금융취약성 간 관계가 지역별 경제·인구 구조와 결합되어 상이한 기여 패턴으로 나타날 수 있음을 시사한다(그림 3, 그림 4).

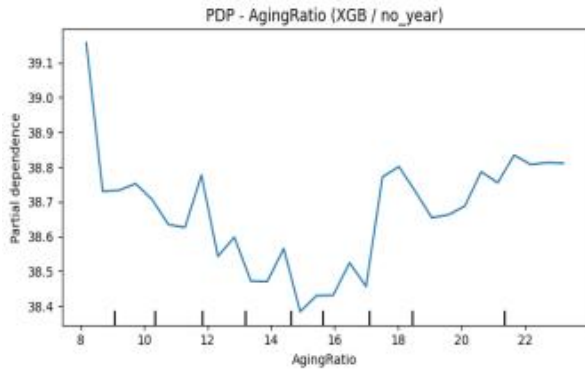


그림 3. 고령화율에 대한 부분의존도 분석 결과
Fig. 3. Partial dependence plot for aging ratio

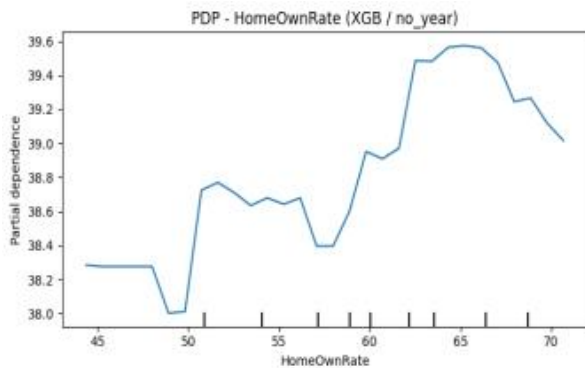


그림 4. 자가 주택 점유에 대한 부분의존도 분석 결과
Fig. 4. Partial dependence plot for homeownership rate

이를 토대로, 금융취약성은 과거 취약성의 경로 의존성과 높은 연관성을 보이며, 동시에 인구·소득·주택구조 변수들이 지역별로 상이한 기여 패턴을 나타내는 구조적 이질성을 가짐을 알 수 있다. 이러한 결과는 이후 군집 분석 및 충격 시나리오 분석에서 지역 구조 유형을 구분하기 위한 해석적 근거로 활용될 수 있다.

4.3 주요 거시 변수의 비선형 효과 분석(PDP)

본 연구에서는 SHAP 분석을 통해 주요 변수의 상대적 중요도를 확인한 후, 각 거시경제 변수가 금융취약성지수(FVI)에 미치는 평균적 영향 구조를

보다 정밀하게 파악하기 위하여 부분의존도(PDP) 분석을 수행하였다. PDP는 특정 변수를 변화시킬 때 다른 변수들의 평균적 효과를 통제된 상태에서 예측값이 어떻게 변화하는지를 나타내는 방법으로, 변수와 종속변수 간의 비선형적 관계 및 임계구간을 확인하는 데 유용하다.

분석 결과, 고령화율은 명확한 비선형적 관계를 보였다. 고령화율이 중간 수준(약 13~16% 구간)일 때 FVI가 상대적으로 낮게 유지되었으나, 고령화가 심화되는 구간(약 18% 이상)에서는 FVI가 다시 상승하는 경향이 관찰되었다. 이는 고령화가 일정 수준을 초과할 경우 금융취약성에 대한 구조적 부담이 확대될 가능성을 시사하며, 고령화의 영향이 단선적 방향성을 갖지 않음을 보여준다.

자가주택비율 또한 비선형적 관계를 나타냈다. 자가비율이 약 60~65% 구간에서 FVI가 상대적으로 높은 수준을 보였으며, 그 이전 및 이후 구간에서는 완화되는 형태가 관찰되었다. 이는 자가주택 구조가 일정 수준 이상 집중될 경우 금융취약성이 확대될 수 있음을 시사하며, 주택 보유 구조가 금융안정성과 복합적으로 연결되어 있음을 보여준다.

실업률 역시 비선형적 관계를 보였다. 낮은 구간에서 FVI가 상대적으로 높은 값을 보인 후, 중간 구간에서는 비교적 안정적인 수준을 유지하고, 일정 수준 이상에서는 오히려 감소하는 양상이 나타났다. 이는 실업률이 금융취약성에 미치는 영향이 단일 방향으로 작용하기보다는, 지역의 산업구조 및 인구 구성과 상호작용하며 복합적으로 작용할 가능성을 시사한다(그림 5).

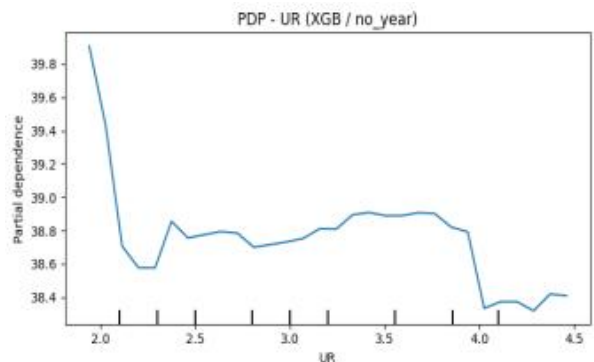


그림 5. 실업률에 대한 부분의존도 분석 결과
Fig. 5. Partial dependence plot for unemployment rate

1인당 GRDP는 특정 수준을 기점으로 FVI가 급격히 하락하는 임계구조를 보였다. 이는 일정 수준 이상의 소득 기반이 확보될 경우 가계의 상환능력 및 지역 재정여건이 개선되면서 금융취약성이 완화될 수 있음을 의미한다. 평균가구원수 역시 특정 구간에서 급격한 변화를 보이며, 가구 구성 구조가 금융취약성에 민감하게 작용할 수 있음을 확인하였다. 임차비율은 중간 구간에서 FVI가 상대적으로 높게 나타났으나, 매우 높은 구간에서는 감소하는 역 U자형 구조를 보였다(그림 6).

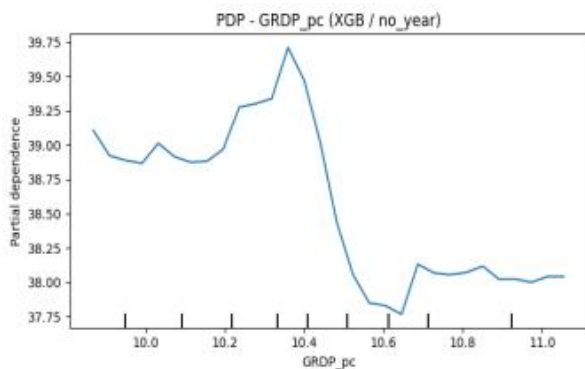


그림 6. 1인당 GRDP에 대한 부분의존도 분석 결과
Fig 6. Partial dependence plot for GRDP per capita

이러한 결과를 통하여, 주요 거시경제 변수들은 금융취약성에 대해 단순 선형 관계가 아닌 비선형적·임계구간적 구조를 갖는 것을 확인할 수 있다. 이는 금융취약성의 결정 요인이 일정 수준을 경계로 영향 강도가 변화하는 구조적 특성을 지니고 있음을 의미하며, 트리 기반 머신러닝 모형이 이러한 복합적 구조를 효과적으로 포착할 수 있음을 뒷받침한다.

4.4 지역 유형화 및 구조적 이질성 분석

본 절에서는 XGBoost 모형으로부터 도출된 지역별 평균 SHAP 값을 활용하여 각 지역이 어떠한 구조적 요인에 의해 금융취약성을 형성하는지를 유형화하였다. 군집 분석에는 지역 $\times$ mean(SHAP) 행렬을 사용하였으며, 변수 간 영향력 크기 차이를 보정하기 위하여 StandardScaler를 적용하였다. 군집 알고리즘은 해석 용이성과 유형 명확성을 고려하여

K-means를 사용하였다.

군집 개수는 실루엣 계수를 기준으로 2~5개를 비교하였다. 그 결과 k=2에서 0.3224로 가장 높은 값을 보였고, k=3에서도 0.3068로 비교적 높은 수준을 유지하였다. 반면 k=4 이상에서는 점수가 급격히 하락하였다. 실루엣 계수 최대값 기준으로는 k=2가 최적이나, 두 집단으로의 구분은 구조 해석이 지나치게 단순화될 가능성이 존재한다. 이에 본 연구는 구조적 유형을 보다 세분화하여 해석하기 위하여 k=3을 최종 군집 수로 설정하였다(표 3).

표 3. 최적 군집 수 결정을 위한 실루엣 계수 분석 결과
Table 3. Silhouette scores for determining the optimal number of clusters

k	Silhouette score
2	0.3224
3	0.3068
4	0.1445
5	0.2866

Cluster 0에는 서울, 경기, 인천, 제주가 포함되었다. 이들 지역은 FVI의 시차향의 영향력이 매우 높게 나타났으며, 실업률의 SHAP 기여도 또한 상대적으로 크게 나타났다. 고령화율 역시 다른 군집 대비 높은 수준을 보였다. 특히 자가주택비율의 영향력이 가장 크게 나타났다는 점은 부동산 구조와 금융취약성 간의 연결성이 강함을 시사한다.

이는 수도권 지역이 경기 변동 및 부동산 시장 구조에 민감하게 반응하는 구조적 특성을 지니고 있음을 의미한다. 과거 취약성이 현재 취약성에 강하게 전이되는 경로의존적 구조가 나타나며, 경기 변동 충격이 취약성으로 전이되는 메커니즘이 존재하는 것으로 해석된다. 따라서 본 군집은 "부동산·경기 민감형 구조"로 명명하였다.

Cluster 1에는 강원, 경남, 경북, 광주, 대구, 대전, 부산, 울산, 전북, 충남, 충북 등 다수의 광역 및 비수도권 지역이 포함되었다. 이 군집은 특정 변수 하나에 과도하게 의존하지 않는 비교적 균형적인 SHAP 구조를 보였다. 고령화율과 실업률의 영향은 상대적으로 낮고, 1인당 GRDP는 높은 수준을 나타내었다.

이는 산업구조와 소득구조가 비교적 안정적으로 작동하며, 취약성이 특정 요인에 집중되지 않는 구조임을 시사한다. 구조적 충격이 발생하더라도 단일 요인에 의해 급격히 취약성이 증폭되는 패턴은 상대적으로 약하다. 따라서 본 군집은 "혼합·중간 안정형 구조"로 정의하였다.

Cluster 2에는 세종과 전남이 포함되었다. 이 군집은 고령화율의 SHAP 기여도가 매우 높게 나타났으며, 평균 가구원수 또한 상대적으로 낮은 값을 보였다. 반면 FVI 시차항의 영향력은 비교적 낮았다.

전남은 고령화가 구조적으로 진행된 지역으로 인구구조 자체가 금융취약성 형성에 중요한 역할을 하는 것으로 해석된다. 세종은 급격한 인구 유입과 구조 변동을 경험한 지역으로, 인구구조 변화에 따른 취약성 변동성이 존재하는 특수 사례로 볼 수 있다. 즉, 이 군집은 경기 요인보다는 인구구조 요인에 의해 취약성이 설명되는 유형이라 할 수 있다. 이에 본 군집을 "인구구조 민감형 구조"로 명명하였다.

SHAP 기반 군집 분석 결과, 지역별 금융취약성은 단순한 수준 차이의 문제가 아니라 어떠한 구조적 요인에 의해 취약성이 형성되는가의 차이로 구분될 수 있음을 확인하였다. 수도권 중심의 부동산·경기 민감형, 다수 비수도권의 혼합·중간 안정형, 그리고 인구구조 중심의 취약성 유형이 구분되었다.

이는 금융취약성 관리 정책이 지역별로 차별화되어야 함을 시사한다. 즉, 수도권은 경기 및 부동산 구조에 대한 거시건전성 관리가 중요하며 고령화가 심화된 지역은 인구구조 대응 정책이 병행되어야 한다.

4.5 외생 충격에 대한 지역별 금융취약성 반응 분석

본 절에서는 SHAP 기반 구조 유형에 따라 지역 금융취약성이 거시경제 충격에 대해 상이한 반응을 보이는지를 분석하였다. 이를 위하여 각 거시경제 변수에 대해 표본 표준편차(1 σ) 수준의 외생적 증가 충격을 가정하고, 다른 변수는 고정된 상태에서 예측 FVI의 평균 변화(Δ FVI)를 산출하였다. 이후 지역별 Δ FVI를 군집 단위로 평균하여 구조 유형별 민감도를 비교하였다.

분석 결과, 구조 유형에 따라 충격 반응의 크기와 방향이 명확히 구분되었다. 먼저 수도권 및 제주로 구성된 Cluster 0(부동산·경기 민감형)은 실업률 충격에 대해 가장 큰 반응을 보였다. 실업률이 1 σ 증가할 경우 해당 군집의 평균 Δ FVI는 -2.48로 나타났으며, 이는 Cluster 1(-0.95) 및 Cluster 2(-0.02)에 비해 절대값 기준 가장 큰 변화이다. 이러한 결과는 SHAP 분석에서 확인된 바와 같이 경기 변수의 기여도가 높고 과거 취약성에 대한 의존도가 강한 구조적 특성을 반영한다. 또한 자가주택비율 충격에 대해서도 +0.54의 증가 반응을 보였으며, 이는 Cluster 1(-0.42)과 상반되는 방향이고 Cluster 2(+0.25)보다 큰 수준이다. 반면 임차비율 충격에는 -0.22의 반응이 나타나 주택 보유 구조 변화가 군집별로 비대칭적으로 작용함을 보여준다. 이러한 결과는 해당 군집이 경기 및 부동산 구조 변화에 구조적으로 민감한 유형임을 실증적으로 뒷받침한다.

이에 비해 강원·영남·충청권 다수 지역으로 구성된 Cluster 1(혼합형·중간 안정형)은 대부분의 거시변수 충격에 대해 상대적으로 완만한 반응을 보였다. 실업률 충격에 대한 Δ FVI는 -0.95로 Cluster 0 대비 약 40% 수준의 절대값에 그쳤으며, 고령화율 충격은 -0.02로 거의 영향이 관찰되지 않았다. 자가주택비율 충격에 대해서는 -0.42로 나타나 Cluster 0(+0.54)과는 방향이 상반되었다. 이는 해당 군집이 특정 단일 변수에 과도하게 의존하지 않는 비교적 균형적인 구조를 가지고 있음을 시사한다. SHAP 구조 분석에서 도출된 "특정 요인 집중도가 낮은 혼합형 구조"라는 해석과도 일관된 결과이다. 즉, 구조적 분산도가 높을수록 외생적 충격에 대한 완충 효과가 나타날 가능성이 있음을 보여준다.

한편 세종과 전남으로 구성된 Cluster 2(인구구조 민감형)는 고령화율 및 가구구조 변수 충격에 대해 상대적으로 높은 반응을 보인 반면, 실업률 충격에는 거의 반응하지 않는 특징을 나타냈다. 평균 가구원수(AVGHHSize) 충격에 대한 Δ FVI는 +0.86으로 세 군집 중 가장 큰 값을 기록하였으며, 이는 Cluster 0(+0.17) 및 Cluster 1(+0.18)에 비해 약 4~5배 높은 수준이다. 반면 실업률 충격은 -0.02에 불과하여 경기 변수에는 거의 민감하지 않은 구조임이

확인되었다. SHAP 분석에서 확인된 바와 같이 이 군집은 고령화율의 기여도가 상대적으로 높고 관성 효과는 낮은 구조를 보였으며, Shock 분석 역시 이러한 구조적 특성이 실제 충격 반응으로 구현됨을 보여준다. 특히 전남의 경우 고령화 심화 지역이라는 특성이 반영된 결과로 해석되며, 세종은 인구 유입과 구조 변화가 빠른 특수 지역이라는 점에서 인구구조 변동성이 금융취약성에 직접적으로 연결될 가능성을 시사한다.

Shock 시나리오 분석을 통해, 본 연구 성과가 실제 거시경제 충격에 대한 반응 간의 연결성을 설명할 수 있음을 알 수 있다. 예컨대, 금융취약성은 동일한 충격 하에서도 지역의 구조적 특성에 따라 상이하게 변화하며, 수도권형은 경기 및 부동산 요인에(실업률 ΔFVI -2.48, 자가주택비율 +0.54), 인구구조 민감형은 가구구조 요인에(평균 가구원수 ΔFVI +0.86) 보다 선택적으로 반응하는 양상이 나타난다. 이는 거시건전성 정책 설계에 있어 전국 단일 대응이 아니라, 지역 구조 유형을 고려한 차별적 접근이 필요함을 시사한다.

V. 결 론

본 연구는 머신러닝 기반 구조 분석을 통해 지역별 금융취약성지수(FVI)의 형성 경로와 거시경제 변수의 영향 구조를 실증적으로 분석하였다. 특히 FVI의 전기값과 지역 거시경제 변수를 동시에 고려한 트리 기반 모형(XGBoost)을 활용하여 단순 예측 정확도 비교를 넘어 구조적 영향 요인을 해석하는데 초점을 두었다.

분석 결과, 금융취약성은 모든 지역에서 동일한 경로로 형성되지 않으며 거시경제 변수에 대한 반응 역시 구조 유형에 따라 상이하게 나타났다. SHAP 분석을 통해 실업률, 고령화율, 평균 가구원수, 주택 점유 구조 등의 변수가 지역별로 서로 다른 기여도를 보였으며, 지역 평균 SHAP 값을 기반으로 수행한 군집 분석에서는 부동산·경기 민감형, 혼합형, 인구구조 민감형 등 구조적 유형이 구분되었다. 또한 1표준편차(1σ) 수준의 외생적 충격을 가정한 Shock 시나리오 분석에서는 동일한 거시경제

충격이더라도 군집별로 ΔFVI 의 크기와 방향이 다르게 나타남을 확인하였다. 이는 금융취약성이 지역의 구조적 특성에 의해 선택적으로 반응함을 의미한다.

이러한 결과는 지역 금융안정 정책 또한 지역별 구조 특성을 반영하여 차별적으로 접근될 필요가 있음을 시사한다. 예를 들어 수도권 중심의 부동산·경기 민감형 구조 지역은 실업률 및 주택시장 변수에 대한 민감도가 높게 나타난 만큼, 부동산 시장 과열 및 경기 변동에 대한 거시건전성 관리가 중요할 수 있다. 반면 인구구조 민감형 지역은 고령화율과 가구구조 변화에 대한 반응이 상대적으로 크게 나타났으므로, 지역 인구구조 변화와 연계된 장기적 금융안정 정책이 병행될 필요가 있다. 또한 혼합형 구조 지역은 특정 변수에 대한 집중도가 상대적으로 낮은 만큼 산업·고용 구조 전반의 안정성을 유지하는 균형적 접근이 요구될 수 있다. 이는 금융취약성이 전국적으로 동일한 방식으로 형성되지 않으며, 지역 구조 유형에 따라 상이한 정책 대응 방향이 필요함을 시사한다.

본 연구의 학술적 기여는 다음과 같다. 첫째, 기존 연구가 주로 예측 성능 비교에 집중된 것과 달리, 머신러닝을 구조 해석 도구로 활용하여 금융취약성의 형성 경로를 실증적으로 제시하였다. 둘째, SHAP 기반 지역 구조 프로파일을 도출하고 이를 군집화 및 Shock 분석과 연결함으로써, 구조 유형과 충격 반응 간의 연계성을 확인하였다. 셋째, 지역 금융안정 정책 설계에 있어 단일 지표 기반 접근의 한계를 보완하고, 구조 유형을 고려한 차별적 접근의 필요성을 제시하였다. 또한, 본 연구는 지역 금융취약성을 구조적 경로의 차이로 접근함으로써 지역 금융안정 연구의 분석 틀을 확장하였다는 점에서 의의를 가지며, 분석 결과를 토대로 데이터 확장과 고도화된 시계열 모형을 적용한 예측 중심 연구의 발전 가능성을 제시한다.

이러한 성과가 있으나, 본 연구 역시 몇 가지 한계를 가진다. 첫째, 본 연구는 17개 지역, 20개 연도(연간 자료)를 기반으로 분석이 이루어져 전체 표본수가 제한적이라는 제약이 있다. 표본 규모가 상대적으로 작은 상황에서는 모델의 일반화 가능성 및 동적 구조 추정에 한계가 발생할 수 있으며, 특히

연간 단위 자료의 특성상 단기 충격에 대한 금융취약성의 동적 반응을 충분히 반영하는 데에는 제약이 존재한다. 또한 학습 데이터와 검증 데이터의 성능 비교 결과 일정 수준의 과적합 가능성이 확인되었으며, 이는 제한된 패널 표본과 구조적 충격에 따른 시계열 변동성의 영향을 반영하는 것으로 판단된다. 이에 본 연구는 모델 복잡도 제어(max_depth, n_estimators, 정규화 설정)와 Panel Cross Validation 기반 강건성 검증을 추가적으로 수행하여 결과의 일반화 가능성을 점검하였다. 또한 지역 - 연도 기반의 균형 패널 구조와 시차 변수를 활용함으로써 금융취약성의 누적적 특성과 지역 간 구조 차이를 최대한 반영하고자 하였다. 향후 월별 또는 분기별 자료 확보를 통해 보다 안정적인 일반화 성능 검증이 가능할 것으로 판단된다.

둘째, 사용된 거시경제 변수 또한 공개 통계 자료의 가용성에 의존하고 있어 산업구조, 금융기관 밀도, 지역 신용공급 특성 등 보다 미시적인 요인을 충분히 반영하지 못하였다.

향후 연구에서는 월별 또는 분기별 단위의 금융-거시 데이터를 추가 확보하여 표본 규모를 확장할 필요가 있다. 데이터의 시계열 해상도를 높일 경우, 보다 정교한 동적 예측 모형 구축이 가능하며 TCN, LSTM 등 딥러닝 기반 시계열 모형을 활용한 금융취약성 예측 연구로의 확장이 기대된다. 특히 장기 의존성과 비선형 동학을 동시에 반영할 수 있는 딥러닝 모형은 지역 금융위험의 조기경보 체계로 발전할 가능성을 가진다.

References

- [1] M. N'Goran, "Measuring Financial Stability and Comparing Financial Risk Monitoring Indicators", *International Economics*, Vol. 185, Art. no. 100681, Mar. 2026. <https://doi.org/10.1016/j.inteco.2025.100681>.
- [2] Bank of Korea, "Compilation of the Financial Vulnerability Index (FVI) and Its Policy Implications", Bank of Korea, Sep. 2021.
- [3] D. Holló, M. Kremer, and M. L. Duca, "CISS: A Composite Indicator of Systemic Stress in the Financial System", *ECB Working Paper Series*, No. 1426, pp. 7-12, Mar. 2012.
- [4] D. Kang, J. Kim, and H. Lee, "Early Warning System in the Korean Financial Markets", *KDI Policy Study* 2005-03, Korea Development Institute, pp. 3-16, Dec. 2005.
- [5] Korea Deposit Insurance Corporation, "Exploring the Financial Conditions Index and Financial Stress Index and Their Usefulness", *Journal of Financial Stability Studies*, Vol. 7, No. 1, pp. 113-134, Jun. 2006.
- [6] R. Cardarelli, S. Elekdag, and S. Lall, "Financial Stress, Downturns, and Recoveries", *IMF Working Paper WP/09/100*, International Monetary Fund, pp. 7-19, May 2009.
- [7] C. Borio and M. Drehmann, "Assessing the Risk of Banking Crises—Revisited", *BIS Quarterly Review*, pp. 29-46, Mar. 2009.
- [8] M. L. Duca and T. A. Peltonen, "Assessing Systemic Risks and Predicting Systemic Events", *Journal of Banking & Finance*, Vol. 37, No. 7, pp. 2183-2195, Jul. 2013. <https://doi.org/10.1016/j.jbankfin.2013.03.010>.
- [9] Financial Supervisory Service, "Prediction and Assessment of Household Debt Default Risk: A Regional Analysis Using Household-Level Data", *Financial Supervisory Research*, Vol. 3, No. 1, pp. 75-101, Jun. 2016.
- [10] L. Breiman, "Random Forests", *Machine Learning*, Vol. 45, No. 1, pp. 5-32, Oct. 2001. <https://doi.org/10.1023/A:1010933404324>.
- [11] T. Chen and C. Guestrin, "XGBoost: A Scalable Tree Boosting System", *Proc. of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, San Francisco, California, USA, pp. 785-794, Aug. 2016. <https://doi.org/10.1145/2939672.2939785>.
- [12] S. Gu, B. Kelly, and D. Xiu, "Empirical Asset Pricing via Machine Learning", *Review of*

Financial Studies, Vol. 33, No. 5, pp. 2223-2273, May 2020. <https://doi.org/10.1093/rfs/hhaa009>.

- [13] J. H. Friedman, "Greedy Function Approximation: A Gradient Boosting Machine", Annals of Statistics, Vol. 29, No. 5, pp. 1189-1232, Oct. 2001. <https://doi.org/10.1214/AOS/1013203451>.
- [14] G. Ke, Q. Meng, T. Finley, T. Wang, W. Chen, W. Ma, Q. Ye, and T.-Y. Liu, "LightGBM: A Highly Efficient Gradient Boosting Decision Tree", Advances in Neural Information Processing Systems, Vol. 30, pp. 1-8, Dec. 2017.
- [15] S. M. Lundberg and S.-I. Lee, "A Unified Approach to Interpreting Model Predictions", Advances in Neural Information Processing Systems, Vol. 30, pp. 1-10, Dec. 2017.
- [16] L. S. Shapley, "A Value for n-Person Games", Contributions to the Theory of Games II, Princeton University Press, pp. 307-317, Jan. 1953.
- [17] H. Li and W. Wu, "Loan Default Predictability with Explainable Machine Learning", Finance Research Letters, Vol. 60, Art. no. 104867, Feb. 2024. <https://doi.org/10.1016/j.frl.2023.104867>.
- [18] M. K. Nallakuruppan, H. Chaturvedi, V. Grover, B. Balusamy, P. Jaraut, J. Bahadur, V. P. Meena, and I. A. Hameed, "Credit Risk Assessment and Financial Decision Support Using Explainable Artificial Intelligence", Risks, Vol. 12, No. 10, Art. no. 164, Oct. 2024. <https://doi.org/10.3390/risks12100164>.
- [19] J. MacQueen, "Some Methods for Classification and Analysis of Multivariate Observations", Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability, pp. 281-297, Jun. 1967.
- [20] S. Lloyd, "Least Squares Quantization in PCM", IEEE Transactions on Information Theory, Vol. 28, No. 2, pp. 129-137, Mar. 1982. <https://doi.org/10.1109/TIT.1982.1056489>.

저자소개

임 해 리 (Haeri Lim)



2024년 3월 ~ 현재 :
서울여자대학교
데이터사이언스학과 학부과정
관심분야 : 머신러닝 기반 예측
분석, 금융 데이터 분석,
공공데이터 기반 의사결정

송 진 아 (Jina Song)



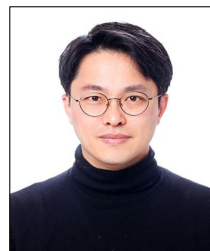
2023년 3월 ~ 현재 :
서울여자대학교
데이터사이언스학과 학부과정
관심분야 : 데이터 기반 서비스
기획, 데이터 시각화, 머신러닝
기반 예측 분석

강 효 연 (HyoYeon Kang)



2024년 3월 ~ 현재 :
서울여자대학교
데이터사이언스학과 학부과정
관심분야 : 빅데이터 활용, 마케팅
성과 분석, 데이터 기반 의사결정

이 종 태 (Jongtae Lee)



1993년 3월 ~ 2001년 2월 :
서울대학교 동물자원학과(학사)
2001년 3월~2007년 2월 :
서울대학교 농경제사회학부
경제학(석사)
2007년 2월~2012년 2월 :
한국과학기술원 경영학(박사)

2015년 9월 ~ 현재 : 서울여자대학교
경영학과/데이터사이언스학과 부교수
관심분야 : IT서비스개발, 정보화전략, 데이터마이닝 기반
전략수립, 공공정보화 정책 방향
분석