

다분류 Support Vector Machine을 이용한 한국 기업의 지능형 기업채권평가모형*

안현철

한국과학기술원 테크노경영대학원 박사후 연구원
(hcahn@kaist.ac.kr)

김경재(교신저자)

동국대학교 경영대학 경영정보학과 조교수
(kjkim@dongguk.edu)

한인구

한국과학기술원 테크노경영대학원 교수
(ighan@kgsu.kaist.ac.kr)

신용등급은 투자자나 채권자 등 다양한 이해관계자들이 특정 기업이나 그 기업에서 발행된 채권에 대한 위험을 평가하는 지표로서, 정교한 등급평가는 개인의 투자위험 뿐만 아니라 금융시장 전체에 영향을 미칠 수 있는 중요한 요소 중 하나이다. 이러한 이유로 지금까지 기업 신용등급평가에 대한 다양한 연구가 진행되어 왔으며, 최근에는 특히 복잡한 재무데이터의 특성을 모형에 보다 잘 반영할 수 있는 것으로 알려진 인공지능기법, 특히 인공신경망의 우수한 예측능력을 활용한 연구가 활발하게 진행되고 있다. 그러나, 인공신경망 기법은 입력자료 분포를 추정하기 위해 다량의 학습데이터가 필요하고, 과도적합문제(overfitting)로 인해 일반화의 어려움이 있을 뿐만 아니라, 지역 최소값(local minima)을 피하기 위한 초기화 작업이 경험에 의존해야 하고, 기본적으로 '암상자 모형'이라서 각 변수의 중요도 등 모형을 해석하기 어렵다는 점 등이 한계점으로 지적되어 왔다. 특히, 기업채권의 등급평가와 같이 다분류 문제의 경우에는 각 등급별 데이터가 희소하여 인공신경망처럼 다량의 학습데이터를 필요로 하는 모형은 구축이 불가능한 경우가 발생할 수 있다.

본 연구에서는 이에 대한 해결방안으로 최근 각광 받고 있는 다분류 support vector machine (SVM)을 채권등급평가에 적용하고자 한다. SVM은 명백한 이론적 근거에 기반하므로 결과 해석이 용이하고, 실제 응용에 있어서 인공신경망 수준의 높은 성과를 내며, 적은 학습자료만으로 신속하게 분류학습을 수행할 수 있다는 장점을 갖고 있다. 또한 기존의 학습 알고리즘은 경험적 위험 최소화 원칙(empirical risk minimization)을 구현하는 것인데 비해, SVM은 구조적 위험 최소화 원칙(structural risk minimization)에 기반하므로 과도적합문제를 어느 정도 피할 수 있다는 장점도 갖고 있다. 본 연구에서는 이 같은 가능성을 확인해 보기 위해, 다분류 SVM을 한국기업의 채권평가 사례에 적용해 보았다. 타 비교모형에 대한 우월성을 검증해 보기 위해, 인공신경망 및 다중판별분석과 그 성과를 비교하였으며, 분석 결과 다분류 SVM이 다른 비교대상들에 비해 통계적으로 유의하게 우수한 성과차이를 보임을 확인할 수 있었다.

주제어: 다분류 SVM, 채권등급평가, 인공신경망, 다중판별분석

1. 서론

기업채권에 대한 정교한 등급평가는 개인의 투자 위험뿐만 아니라 금융시장 전체에 영향을 미칠 수

있는 중요한 요소 중 하나이다. 따라서 채권등급평가는 채권자나 주주는 물론 종업원, 고객, 소비자 등 모든 경제주체들에게 있어서 중요한 관심사라고 할 수 있다.

상당히 많은 연구들이 정교한 채권등급평가를 위

한 모형을 개발하여 왔으나 주로 통계적인 기법을 활용한 연구가 많았으며, 최근에는 복잡한 재무데이터의 특성을 모형에 보다 잘 반영할 수 있는 것으로 알려진 인공지능기법을 이용한 연구들도 진행되고 있다. 특히, 인공지능기법을 이용한 연구들 중에는 인공신경망의 우수한 예측능력을 활용한 연구가 활발하게 진행되고 있다. 그러나, 인공신경망을 사용할 경우 입력자료의 분포를 추정하기 위해 다량의 학습데이터가 필요하고, 과도적합문제(overfitting)로 인해 일반화의 어려움이 있을 뿐만 아니라, 지역 최소값(local minima)을 피하기 위한 초기화 작업이 경험에 의존해야 하고, 기본적으로 '암상자 모형'이라서 각 변수의 중요도 등 모형을 해석하기 어렵다는 점 등이 한계점으로 지적되어 왔다. 특히, 기업채권의 등급평가와 같이 다분류 문제의 경우에는 각 등급별 데이터가 희소하여 인공신경망처럼 다량의 학습데이터를 필요로 하는 모형은 구축이 불가능한 경우가 발생할 수 있다.

본 연구에서는 이에 대한 해결방안으로 최근 각광 받고 있는 support vector machine(SVM)을 채권등급평가에 적용하고자 한다. SVM은 1995년 Vapnik에 의해 제안된 학습이론으로 분류문제를 해결하기 위한 최적의 분리 경계면(hyperplane)이라는 개념을 사용한다. SVM이 주목 받는 이유는 첫째, 명백한 이론적 근거에 기반하므로 결과 해석이 용이하고, 둘째, 실제 응용에 있어서 인공신경망 수준의 높은 성과를 내고, 셋째, 적은 학습자료만으로 신속하게 분류학습을 수행할 수 있기 때문이다. 또한 SVM은 기존의 학습 알고리즘이 경험적 위험 최소화 원칙(empirical risk minimization)을 구현하는 것인데 비해 구조적 위험 최소화 원칙(structural risk minimization)에 기반하므로 과도적합문제를 어느 정도 피할 수 있다. 기존의

연구에서는 SVM을 기업부도예측이나 주가지수의 방향성 예측 등 이분류 문제에 주로 적용하여 왔으나 본 연구에서는 다분류 문제를 해결할 수 있는 다분류 SVM(multi-class SVM)을 활용하여 다분류 문제의 대표적인 예인 기업채권등급평가문제에 적용하고자 한다.

본 연구는 총 5장으로 구성하였다. 1장에서는 연구의 의의와 목적을 정의하고, 2장에서는 기존의 채권등급평가모형에 관련된 선행연구와 SVM에 대해 알아보도록 한다. 이어서 3장에서는 채권등급평가를 위한 SVM모형을 제안하고 다중판별분석, 인공신경망 등의 기존 방법론들과 비교분석을 한다. 그리고, 4장에서는 3장의 연구 결과를 분석하고, 5장에서는 본 연구의 시사점과 결론을 제시한다.

II. 선행연구

채권등급평가에 관해서는 많은 연구자들에 의해 이미 상당한 연구가 이루어졌다. 본 장에서는 기존 연구들을 방법론에 따라 통계적 기법을 사용한 모형과 인공지능기법을 사용한 모형으로 나누어 살펴보고, 본 연구에서 제안하고자 하는 support vector machine에 대해 살펴 본다.

2.1 통계적 기법을 이용한 채권등급평가모형

통계적 모형을 이용한 채권등급평가에 관한 실증분석을 수행한 연구는 Fisher(1959)의 연구를 시작으로 볼 수 있다. Fisher는 OLS(ordinary least square)를 활용하여 채권의 위험에 대한 보상(risk premium)의 분산을 설명하려 하였으며, 이 연구와

같이 OLS를 채권등급평가에 활용한 연구로는 Pogue & Solofsky(1969), West(1991) 등이 있다. 채권등급평가와 같이 다분류 문제에 적합한 통계적 기법 중 대표적인 것은 다중판별분석(multiple discriminant analysis; MDA)이며, 이를 활용한 연구로는 Pinches & Mingo(1973; 1975)가 있다. 다중판별분석은 독립변수군과 종속변수 사이의 선형판별함수를 추정하는 기법으로 비교적 높은 판별력과 설명력을 가지고 있다는 장점이 있다. 통계적 기법을 채권등급평가에 활용한 기타 연구로는 로지스틱 회귀분석을 활용한 Ederington(1985)과 프로빗(probit)을 활용한 Gentry et al.(1988), Jackson & Boyd(1988) 등이 있다. 로지스틱 회귀분석과 프로빗은 기본적으로 이분류 문제에 주로 많이 활용되는 통계적 기법으로 선행연구에서는 보다 활용도가 높은 다중판별분석이 주로 활용되었으며, 이후의 연구에서도 통계적 기법을 활용한 연구가 진행되었으나 주로 인공지능기법과의 성과 비교를 위해 활용된 경우가 많았다. 이는 인공지능기법이 통계적 기법에 비해 재무데이터를 처리하는데 있어서 여러 가지 장점이 있기 때문이며, 이에 대해 다음 절에서 설명한다.

2.2 인공지능기법을 이용한 채권등급평가모형

1980년대 후반부터는 인공지능망, 귀납적 학습방법, 사례기반추론, 유전자 알고리즘 등 인공지능기법이 재무예측 분야에 응용되기 시작하였다. 이중 가장 대표적인 인공지능기법은 인공지능망이며, 이를 채권등급평가에 적용한 초기 연구는 Dutta & Shekhar(1988)가 있다. 이 연구에서는 채권등급을 두 개의 범주로 분류하는 문제에 오류역전과 인공지능망을 적용하여 인공지능망의 채권등급평가

능성을 평가하였다. 통계적 기법과의 본격적인 비교를 수행한 연구로는 Kim(1993)이 있으며, 선형회귀분석, 판별분석 등의 전통적인 통계적 기법과의 비교를 하였고 인공지능망의 성과가 가장 우월함을 보였다.

이 밖에 채권등급평가에서 다중판별분석과 인공지능망의 성과를 비교한 연구로는 Lee et al.(1996), Kwon et al.(1997), Maher & Sen(1997) 등이 있고 이 중 Kwon et al.(1997)은 다중판별분석과 OPP(ordinal pairwise partitioning)를 활용한 인공지능망 모형의 성과를 비교하였다. 세 연구의 결과는 인공지능망의 성과가 다중판별분석보다 우월하다는 것이었다. Chaveesuk et al.(1999)는 전통적인 통계적 기법의 하나인 로지스틱 회귀분석과 인공지능기법인 RBF, LVQ 등을 인공지능망 모형과 비교하였는데 이 연구에서도 역시 인공지능망의 채권등급평가 결과가 가장 우수함을 보였다. 이들 연구들을 통해 일반적으로 인공지능망 모형은 전통적인 통계적 기법에 기반한 모형보다 채권등급평가에 있어서 우월한 성과를 보이는 것으로 생각할 수 있다.

이후 연구들은 인공지능망 이외에 다른 인공지능기법들의 결합모형도 전통적인 모형들에 비해 채권등급평가에 있어서 유용함을 보였다. Shin & Han(1999)은 유전자 알고리즘과 사례기반추론의 결합모형을 이용하여 채권등급평가를 하고 그 결과를 다중판별분석, 의사결정나무, 일반적인 사례기반추론과 전문가의 의견을 반영한 사례기반추론 등의 결과와 비교하여 제안한 모형의 성과가 가장 우수하였다고 주장하였다. Shin & Han(2001)은 의사결정나무와 사례기반추론이 결합된 모형의 채권등급평가 결과를 다중판별분석, ID3, 일반적인 사례기반추론의 결과와 비교하였으며, 그 결과 결합모형이

다중판별분석 등 다른 모형보다 우월함을 제시했다.

이처럼 많은 연구들을 통해 입증된 인공신경망의 우수한 예측정확성에도 불구하고, 인공신경망의 결과는 설명력이 부족하여 예측결과의 원인을 설명하기 어렵다는 한계점과 함께 소량의 자료를 대상으로 실험을 하였을 때 발생할 수 있는 일반화의 어려움도 한계점으로 제기되었다(Jo & Han, 1996). 특히, 이 문제는 채권등급평가와 같은 다분류 문제의 경우에 각 등급별로 데이터가 매우 희소한 경우가 발생할 수 있으므로 중용한 한계점이 될 수 있다. 또한, 모형을 구축함에 있어 과도적합의 문제와 신경망 구조의 설계에 대해 많은 시간과 노력이 필요하다는 단점도 제기되었다(Altman et al., 1994).

2.3 Support Vector Machine (SVM)

SVM은 통계학자인 Vapnik에 의해 개발된 분류 기법으로, 입력공간과 관련된 비선형문제를 고차원의 특징공간에서의 선형문제로 대응시켜 나타내기 때문에 수학적으로 분석하는 것이 수월하다(Vapnik, 1995, Hearst et al., 1998). 또한, SVM은 조정해야 할 파라미터의 수가 많지 않아 비교적 간단하게 학습에 영향을 미치는 요인들을 규명할 수 있다. 그리고 구조적위험을 최소화함으로써 과도적합문제에서 벗어날 수 있으며, 볼록함수를 최소화하는 학습을 진행하기 때문에 전역최적해(globally optimal solution)를 구할 수 있다는 점에서 인공신경망보다 우월한 기계학습기법으로 주목 받고 있다.

최근 몇 년간 SVM을 사용한 다양한 연구가 진행되었는데 SVM은 문서분류, 영상인식, 문자인식 등에서 뛰어난 일반화 성능을 보여주었다(Joachims, 1998, Osuna et al., 1997). 또한, SVM을 재무 분야에 적용한 연구도 있는데, 주로 시계열 예측 및

분류에 관한 것이었다(Tay & Cao, 2002, Kim, 2003). 이 연구들에서 SVM은 일반화에 있어서 인공신경망이나 판별분석 등의 다른 분류기법들과 비교하여 비슷하거나 더 우수한 성능을 나타낸 것으로 보고되었다. 본 연구에서는 이러한 연구 배경을 토대로 시계열 분류문제와는 다른 재무적 특성을 지닌 기업채권등급평가에 SVM을 적용하여 보기로 한다. 기업채권은 일반적으로 여러 개의 등급으로 나뉘어 지므로 전형적인 다분류 문제이나 고전적 SVM은 원래 이분류를 위한 알고리즘으로 고안되었다. 따라서 본 절에서는 먼저 이분류를 위한 일반적인 SVM의 원리에 대해 설명하고 다분류를 위한 SVM에 대해서는 이후에 차이점만 간단히 설명하고자 한다.

SVM에서는 모형구축용 데이터들을 서로 다른 두 개의 클래스로 분류할 때 분류의 기준이 되는 분리 경계면(hyperplane)을 학습 알고리즘을 이용하여 찾는다(이수용, 이일병, 2002). 따라서, SVM은 입력벡터 x 를 고차원의 특징공간(high-dimensional feature space)으로 사상(mapping)시킨 후 두 클래스 사이의 마진(margin)을 최대화시키는 분리 경계면을 찾는 것을 목적으로 한다. 이러한 최대마진 분리 경계면(maximum margin hyperplane)은 두 클래스 사이의 거리를 최대로 분리시킨다. 이 때 최대마진 분리 경계면에 가장 근접한 모형구축용 데이터를 서포트 벡터(support vector)라고 부른다. 선형분리문제에서, 독립변수가 3개인 경우 분리 경계면은 식(1)과 같다.

$$y = w_0 + w_1x_1 + w_2x_2 + w_3x_3 \quad (1)$$

여기서 y 는 출력값이고, x_i 는 변수값, 그리고 4개의 w_i 는 학습 알고리즘에 의해 학습된 가중치이다. 상기 식에서 가중치 w_i 는 분리 경계면을 결정

하는 파라미터이다. 이 때 최대마진 분리 경계면은 서포트 벡터를 사용해서 식(2)와 같이 나타낼 수 있다.

$$y = b + \sum \alpha_i y_i x(i) \cdot x \quad (2)$$

여기서, y_i 는 모형구축용 데이터 $x(i)$ 의 분류값이고, $\cdot$ 는 내적(dot product)이다. 벡터 x 는 모형검증용 데이터를 나타내고, 벡터 $x(i)$ 는 서포트 벡터(최대마진 분리 경계면에 가장 근접한 모형구축용 데이터)를 나타낸다. 이 식에서, b 와 α_i 는 분리 경계면을 결정하는 파라미터이다. 서포트 벡터를 찾아내고, 파라미터 b 와 α_i 를 결정하는 것은 선형적으로 제약된 이차계획문제(linearly constrained quadratic programming)를 푸는 것과 같다.

앞에서 언급한 바와 같이, SVM은 입력변수를 고차원의 특징 공간으로 이동시킴으로써 비선형 분류 문제를 선형모형으로 근사시킨다. 비선형 분류문제에서 사용될 식(2)의 고차원 형태는 식(3)과 같이 나타낼 수 있다.

$$y = b + \sum \alpha_i y_i K(x(i), x) \quad (3)$$

식(3)에서 함수 $K(x(i), x)$ 는 커널함수(kernel function)라고 정의된다. 커널함수는 원래 데이터를 고차원 공간으로 사상시킴으로써 특징공간 내에 선형으로 분리가능한 입력 데이터셋을 만든다. 이 때 사용될 수 있는 커널함수는 여러 가지가 있으며 어떤 커널함수를 선택하는 것이 바람직한가는 문제에 따라 상이하고, 이는 SVM을 적용하는 데 있어서 가장 중요한 요소 중의 하나이다. 일반적으로 많이 사용되는 커널함수로는 선형함수(linear function)

와 다항식 함수(polynomial function), 그리고 가우시안 RBF 함수(Gaussian radial basis function)를 들 수 있으며, 각 함수식은 아래와 같다.

$$\text{선형 함수: } K(x, y) = xy \quad (4)$$

$$\text{가우시안 RBF: } K(x, y) = \exp\left(-\frac{1}{\delta^2}(x - y)^2\right) \quad (5)$$

$$\text{다항식 함수: } K(x, y) = (xy + 1)^d \quad (6)$$

여기서 d 는 다항식 함수의 차수이고, δ^2 은 가우시안 RBF 함수의 대역폭이다.

분리가능한 문제에 있어서 상기 식의 계수 α_i 의 하한은 0이다. 분리가 불가능한 문제에서 SVM은 계수 α_i 의 하한 이외에 상한 C 를 추가함으로써 일반화된 결과를 얻을 수 있다(Kim, 2003).

이상은 고전적인 이분류를 위한 SVM의 원리를 설명하였다. 본 연구의 응용주제인 채권등급평가는 일반적으로 3개 이상의 등급을 갖게 되므로 고전적인 SVM을 사용할 수 없다. 따라서 이분류를 위한 고전적인 SVM을 확장한 다분류 SVM(multi-class SVM)을 사용한다. Hsu & Lin(2002)에 의하면 다분류 SVM을 구성하는 방법은 두 가지로 나누어 볼 수 있는데, 하나는 이분류 SVM을 여러 개 결합하는 방법이고, 다른 하나는 모든 class를 한번에 고려하여 하나의 최적화 문제로 해결하는 방법이다. 우선 전자의 경우, 이분류의 SVM을 다수 개 결합하는 방법에는 분류할 등급의 개수만큼의 SVM모형을 구축해, 각 등급에 해당하는 그룹과 해당하지 않는 그룹을 판별하는 형태의 (1) one-against-all 방법과 분류할 전체 등급에 대해 구성될 수 있는 모든 쌍(pair)별로 독립된 SVM모형을 구축하는 (2) one-against-one 방법의 2가지가 있다. 예를 들어, A, B, C의 3가지 등급이 있다고

할 경우, (1)에서는 A vs. not A, B vs. not B, C vs. not C의 3가지 모형을 구축하게 되며, (2)에서는 A vs. B, A vs. C, B vs. C의 3가지 모형을 구축하게 되는 것이다. 이 때, 분류할 등급의 개수를 k 라 할 때, (1)의 방법은 k 개의 모형을 구축해야 하는 반면, (2)의 방법은 $kC_2 = \frac{k(k-1)}{2}$ 개 만큼의 모형을 구축해야 한다. 따라서, k 가 클 경우, (2)보다는 (1)의 방법이 더 적용이 용이하게 된다.

한편, 모든 class를 한번에 고려하여 하나의 최적화 문제로 해결하는 방법은 기본적으로 앞서 설명한 이분류 SVM 기본 모형과 유사하다. 다만, 최종 추정된 분리경계면을 k 개의 등급에 의해 분류할 수 있어야 하므로, 이차계획문제에 포함된 선형제약식이 k 개 만큼 늘어나게 된다는 차이가 있다. 때문에, 이차계획문제의 해를 찾는 데 있어서 이러한 다분류 SVM 모형은 일반적인 이분류 SVM 모형보다 훨씬 복잡한 분해방식을 요구하게 되는데, Hsu & Lin (2002)는 이러한 문제에 대한 대안으로 이차계획문제의 목적식에 특정한 항($\sum_{m=1}^k b_m^2$)을 임의로 추가해 분해를 용이하게 하는 방법과 Cramer와 Singer가 제안한 여유변수(slack variable)를 활용한 방법 등을 제시하고 있다.¹⁾

이러한 다분류 SVM은 지금까지 다양한 분야에 적용되어 왔다. Foody & Marther(2004)는 '이미지 분류(image classification)'에 있어, 다분류 SVM이 인공신경망이나 의사결정나무와 같은 다른 인공지능기법보다 훨씬 우수한 성과를 보임을 제시하였고, Ramaswamy et al.(2001)과 Statnikov et al.(2005)은 생물정보학(bioinformatics) 분야에 적용하여, 유전자배열을 이용한 암 예측에 있어

서 다분류 SVM이 사례기반추론이나 인공신경망 등 다른 비SVM 기법들에 비해 압도적으로 우수한 예측 성과를 보임을 제시하였다. 본 연구에서 다루고 있는 신용등급평가에 다분류 SVM을 적용한 연구로는 Huang et al.(2004)의 연구가 있다. Huang et al.(2004)은 대만과 미국 기업들의 신용등급 평가에 다분류 SVM과 인공신경망, 로지스틱 회귀분석을 적용해, 다분류 SVM이 로지스틱 회귀분석보다는 우수하고, 인공신경망 기법과 거의 비슷한 수준의 예측력을 보임을 제시하였다. 또한 예측에 사용된 각 입력변수의 상대적 중요도를 도출하여, 이것이 대만 시장과 미국 시장 사이에 서로 차이가 있음도 함께 제시하였다.

이러한 Huang et al.(2004)의 연구는 신용등급 평가에 처음으로 다분류 SVM을 적용했다는 측면에서 그 연구의 의의가 있다고 할 수 있으나, 동시에 다음과 같은 두 가지 한계점을 갖고 있다. 첫째, 저자들도 밝히고 있듯이, 이 연구는 대만기업 74개와 미국기업 265개를 대상으로 한 연구이므로, 우선 연구의 대상이 된 시장이 대만과 미국으로 한정되어 있고, 동시에 연구의 대상이 된 표본이 인공지능기법을 적용하기에는 너무 적은 수라는 한계가 있다. 둘째로, 이 연구에서는 다분류 SVM이 인공신경망보다 전반적으로 우수하거나 유사한 성과를 보인다고 제시하고 있으나, 그 차이가 과연 통계적으로 유의한 차이인지를 검증하는 과정도 생략되어 있다. 때문에 본 연구에서는 다분류 SVM을 한국 내 기업들에 대한 신용등급예측에 적용하되, 그 분석 대상을 1000개 이상의 표본으로 하여, 구축된 모형이 충분한 자료로부터 패턴을 추출할 수 있게끔 한다. 또한, 예측결과의 차이가 통계적으로 유의한지를 살

1) 기술적인 설명에는 많은 수식이 동반되므로 관심 있는 독자는 Hsu & Lin (2002)을 참고하기 바람.

펴보기 위해, McNemar 테스트 등을 적용해 그 결과를 살펴보고자 한다.

III. 데이터와 실험설계

3.1 데이터 수집 및 변수선정

다분류 SVM이 경영/경제와 관련한 다분류 문제에 얼마나 우수한 문제해결 능력을 갖는지 검증해 보기 위해, 본 연구에서는 다분류 SVM을 비롯한 통계 및 인공지능 기법들을 국내 기업들의 신용등급 예측에 적용해 보았다. 대상이 된 데이터는 제조업에 종사하며 KOSPI에 상장되어 있거나 KOSDAQ에 등록되어 있는 1295개의 기업을 표본으로 하여 이들의 회사채 신용등급을 예측의 목적변수로 설정하고 실험을 진행하였다. 이들 기업의 신용등급은 2002년에 발표된 국내 N 신용정보기관의 공시 자료를 이용하였고, 실제 해당 기업들의 재무제표 자료는 한국상장회사협의회에서 제공하는 데이터베이스에서 추출하였다. 일반적으로 기업의 신용등급은 크게 A1, A2, A3, B, C의 5등급으로 나뉘고 있는

데, 본 연구에서는 편의상 A1은 1, A2는 2, A3는 3, B와 C는 4로 표기하였다. B와 C 등급을 하나의 등급으로 묶은 이유는 C등급에 해당하는 사례가 거의 없었으며, 일반적으로 신용평가회사들이 신용등급을 부여할 때 대부분의 기업들에 대해 하한을 B등급으로 부여하는 관행을 갖고 있기 때문에 'B등급 이하'는 그 자체로 투자부적격채권(junk bond)의 의미로 해석할 수 있기 때문이다. 모형의 구축을 위한 자료의 구분과 관련해서는 학습용(training) 데이터로 각 등급의 80%에 해당하는 데이터를 그리고 나머지 20%를 검증용(validation) 데이터로 사용하였다. 한편, 인공지능망과 같은 인공지능 모형은 전술한 바와 같이 각 등급별 데이터가 희소한 경우에 과도적합 등의 문제를 발생시킬 수 있으므로 비교를 위해서는 등급별 데이터 개수를 증가시켜야만 하기에 5-fold cross validation을 실시하였다. 실험 데이터의 최종적인 채권등급 분포는 <표 1>과 같다.

일반적으로 채권등급평가에 적용되는 요인변수는 크게 재무제표로부터 추출되는 재무변수와 기업의 형태, 산업 및 환경요인 등의 비재무변수가 있다. 본 연구에서는 이 중 계량적 분석에 적합한 재무변수들을 활용해 모형을 구축하는데, 기존 연구에서

<표 1> 실험용 데이터의 등급별 분포 현황

실제 등급	부여 등급	1-fold			5-fold			백분율
		학습용	검증용	합계	학습용	검증용	합계	
A1	1	85	21	106	425	105	530	8.19%
A2	2	442	110	552	2210	550	2760	42.63%
A3	3	266	66	332	1330	330	1660	25.64%
B	4	244	61	305	1220	305	1525	23.55%
C								
합 계		1037	258	1295	5185	1290	6475	100.00%

채권등급에 유의한 영향을 주는 것으로 제시되어 온 총 39개의 재무변수들을 후보 변수로 하여, 분석을 진행한다(한인구 등, 1995; 이견창 등, 1996; Shin & Han, 1999; 박기남 등, 2000; Huang & Tseng, 2004 참고). 후보 변수들은 규모지표 11개, 수익성 지표 13개, 안정성 지표 8개, 현금흐름 지표 4개, 생산성 지표 3개로 구성되는데, 그 내역은 다음의 <표 2>와 같다.

이상의 39개 후보변수 중에서, 우선 독립 T-test를 적용해 기업신용등급과 95% 신뢰수준 하에서 유의한 것으로 밝혀진 총 36개의 변수를 1차로 선정한다. 그러나, 이 변수들도 모형을 구축하기에는 매우 많은 수이므로, 본 연구에서는 순차적 다차원 판별분석(Stepwise MDA)을 통해 최종적으로 선택된 총 14개의 변수들만 분석에 활용한다.

3.2 SVM 모형구축

앞의 문헌연구에서 살펴본 바와 같이 SVM 모형의 구축에 있어서 어떤 커널함수를 사용하느냐는 모형의 성과를 결정하는 가장 중요한 문제 중의 하나이다. 본 연구에서는 SVM의 커널함수로서 가장 널리 사용되는 3가지 함수인 선형 커널과 다항식 커널, 그리고 가우시안 RBF를 사용한다. Tay & Cao(2001)에 의하면 SVM의 성능에 있어서 커널함수의 상한 C 와 커널 파라미터 δ^2 , d 가 중요한 역할을 한다고 보고되었다. 따라서 적절한 상한 C 와 커널 파라미터 δ^2 , d 를 선정하기 위해, 선행연구에서 SVM의 파라미터에 대해 제시된 일반적인 가이드를 따라 본 연구에서도 일정 범위 내 다양한 값을 대입하여 다양한 모형을 생성시켰다.

SVM 모형의 구축은 다분류 기능을 지원하는 공개 소프트웨어인 BSVM(<http://www.csie.ntu.edu.tw/~cjlin/bsvm/>)을 사용하였다. BSVM은 효율적이면서도 효과적인 집단 분류를 위해 총 3가지 SVM 문제해결 방식을 선택할 수 있도록 제공하고 있는데, 이는 각각 (1) bound-constrained multi-class support vector classification (SVC)을 이용하는 방법, (2) multi-class SVC from solving a bound-constrained problem을 이용하는 방법, 그리고 (3) multi-class SVC from Crammer and Singer을 이용하는 방법이다. 이 중, 첫 번째 bound-constrained multi-class support vector classification (SVC) 방법은 전체 등급에 대해 구성될 수 있는 모든 쌍별로 이분류 SVM모형을 구축해, 신용등급을 구분 하는 one-against-one 방법이며, (2) multi-class SVC from solving a bound-constrained problem 와 (3) multi-class SVC from Crammer and Singer는 다분류를 할 수 있는 하나의 SVM모형을 구축해, 이 모형을 기반으로 다분류를 수행하는 방법이다. 단 이 때, (2)와 (3)의 차이점은 (2)는 다분류 SVM 모형의 해를 찾는 과정에 있어서, Weston & Watkins(1999)가 제시한 일반적인 분해 기법(decomposition method)을 적용하는 반면, (3)은 Cramer & Singer(2000)가 제안한 새로운 분해 기법을 이용해 해를 찾는다라는 점이다. 이상의 세 가지 기법들은 각각 장, 단점이 동시에 존재하기 때문에, 본 연구에서는 이 3가지 경우를 모두 실험하여, 가장 우수한 실험결과를 선택하는 형태로 연구를 진행한다.

edu.tw/~cjlin/bsvm/)을 사용하였다. BSVM은 효율적이면서도 효과적인 집단 분류를 위해 총 3가지 SVM 문제해결 방식을 선택할 수 있도록 제공하고 있는데, 이는 각각 (1) bound-constrained multi-class support vector classification (SVC)을 이용하는 방법, (2) multi-class SVC from solving a bound-constrained problem을 이용하는 방법, 그리고 (3) multi-class SVC from Crammer and Singer을 이용하는 방법이다. 이 중, 첫 번째 bound-constrained multi-class support vector classification (SVC) 방법은 전체 등급에 대해 구성될 수 있는 모든 쌍별로 이분류 SVM모형을 구축해, 신용등급을 구분 하는 one-against-one 방법이며, (2) multi-class SVC from solving a bound-constrained problem 와 (3) multi-class SVC from Crammer and Singer는 다분류를 할 수 있는 하나의 SVM모형을 구축해, 이 모형을 기반으로 다분류를 수행하는 방법이다. 단 이 때, (2)와 (3)의 차이점은 (2)는 다분류 SVM 모형의 해를 찾는 과정에 있어서, Weston & Watkins(1999)가 제시한 일반적인 분해 기법(decomposition method)을 적용하는 반면, (3)은 Cramer & Singer(2000)가 제안한 새로운 분해 기법을 이용해 해를 찾는다라는 점이다. 이상의 세 가지 기법들은 각각 장, 단점이 동시에 존재하기 때문에, 본 연구에서는 이 3가지 경우를 모두 실험하여, 가장 우수한 실험결과를 선택하는 형태로 연구를 진행한다.

3.3 인공신경망 모형구축

본 연구에서 SVM에 대한 비교대상으로서 인공신경망과 MDA를 수행한다. 각 기법들에 대한 설계는

〈표 2〉 실험에 사용된 후보변수 목록

변수 구분	후보변수
규모 지표	총자산 유형자산 고정자산 자기자본 매출액 부가가치 총부채 감가상각비 영업이익 당기순이익 업력
수익성 지표	주당순이익 유보액대비율 금융비용부담률 금융비용대부채비율 금융비용대총비용비율 감가상각비대총비용비율 이자보상배율 총자본순이익률 자기자본순이익률 자본금영업이익률 자본금경상이익률 매출액총이익률 매출액경상이익률
안정성 지표	고정자산구성비율 차입금의존도 자기자본구성비율 재고자산대유동자산비율 단기차입금대총차입금비율 고정부채대자본금비율 부채비율 (자본+고정부채)/고정자산
현금흐름 지표	현금흐름대부채비율 현금흐름대고정부채비율 현금흐름대총자본비율 (영업활동으로인한현금흐름-현금배당)/(고정자산+운전자본)
생산성 지표	1인당매출액 총자본사업이익률 총자본투자효율

일반적으로 알려진 범위 내에서 가장 우수한 성능을 나타내는 방향으로 진행한다.

전술한 바와 같이 인공신경망은 부도예측 분야에서 가장 많이 사용되어 온 기법이므로 본 연구에서는 인공신경망의 일반적인 작동원리에 관해서는 그 설명을 생략하기로 한다. 한편, 인공신경망 모형의 설계에 관해서는 변수선정과정, 신경망 구조 설계 등에 있어서 일반적인 원칙이 없다. 따라서 본 연구에서도 기존 연구에서 사용된 방법과 같이 다양한 실험조건으로 실험하고 가장 좋은 신경망 구조를 선택한다.

일반적으로 인공신경망의 성능에 영향을 미치는 요인으로서 은닉층의 수, 은닉층의 노드 수, 학습횟수 등이 알려져 있다. Hornik(1991)에 따르면 은닉층의 수는 하나만으로도 분류문제를 포함한 대부분의 문제에서 만족할만한 결과를 얻을 수 있다고 한다. 따라서, 본 연구에서도 은닉층이 하나인 3층 퍼셉트론을 사용하였다. 은닉층의 노드 수는 경험적으로 입력노드 수와 출력노드 수의 합을 n 이라 할 때 $n/2$, n , $2n$ 을 많이 사용하지만, 모든 경우에 적합하다고 할 수는 없다. 은닉노드의 수는 인공신경망 구조를 설계하는데 있어서 중요한 요소이며 그 결정은 데이터 의존적인 경우가 많다. 학습용 데이터를 분류하는 경우에는 은닉노드의 수가 많을수록 바람직하지만, 검증용 데이터에서는 은닉노드의 수가 많은 것이 반드시 바람직한 것은 아니다 (Patuwo et al., 1993). 따라서 은닉노드의 수를 많게 하거나 적게 하는 두 가지 방법 모두 득실이 있으므로 본 연구에서는 은닉노드의 수를 7, 14, 21, 28로 구분하여 실험해 본다.

일반적으로 인공신경망의 경우, 과도적합을 막기 위해, 학습용 데이터를 다시 학습용과 테스트용의 2가지로 구분하여, 모형이 학습용 데이터에 너무 과

도적합하는 것이 아닌지 테스트용 데이터를 이용해 판별하게 된다. 이에 본 연구에서는 인공신경망 실험의 경우, 학습용 데이터를 각각 3:1의 비율로, 실제 학습에 사용되는 데이터와 과도적합을 방지하기 위한 테스트용 데이터를 구분한다. 아울러, 학습횟수는 너무 적으면 학습이 제대로 이루어지지 않고, 너무 많아도 학습용 데이터에 과도적합되어 검증용 데이터의 예측력이 떨어지는 경우가 많다. 따라서, 본 연구에서는 학습이 적당히 이루어지도록 실제 학습용으로 활용된 데이터 개수 779건의 약 50배에 해당되는 38,950회가 지나면 학습이 멈추도록 한다. 편의를 위해 인공신경망의 학습률(learning rate)은 0.1, 모멘텀은 0.1로 고정한다.

3.4 다변량판별분석(MDA) 모형구축

다변량판별분석은 로지스틱 회귀분석과 마찬가지로 등간척도나 비율척도로 측정된 독립변수와 명목척도인 종속변수를 이용한 분석기법으로 선형적으로 정의된 두 개 이상의 집단들을 가장 잘 판별할 수 있는 둘 이상의 독립변수의 선형조합을 찾아내는 과정을 포함한다. 다변량판별함수는 식(7)과 같은 형태이다.

$$Z = W_1X_1 + W_2X_2 + W_3X_3 + \dots + W_nX_n \quad (7)$$

여기서 Z 는 판별점수이고, W 는 판별계수이고, X 는 독립변수를 말한다. 다변량판별분석을 위해 본 연구에서는 SPSS 소프트웨어를 사용하여 수행한다.

IV. 연구결과

본 연구에서는 SVM의 실험결과를 각 커널함수와 파라미터에 따라 정리해보고, 추가적으로 인공신경망, MDA의 실험결과와 비교해보고자 한다. 전술한 바와 같이 본 연구에서 사용하는 커널함수는 가우시안 RBF와 다항식 및 선형 커널이다. 따라서 본 연구에서는 상한 C 와 커널 파라미터를 변경하면서 실험을 진행한다.

선형커널의 경우, 상한 C 만 변경하면 되기 때문에, 이 값을 다양하게 변경시켜가면서 실험을 진행한다. 반면 다항식 커널의 경우에는 C 와 동시에 차수를 의미하는 d 도 함께 조정한다. 가우시안 RBF에서는 C 이외에 커널 파라미터로 δ^2 을 고려해야 한다. Tay & Cao(2001)에 따르면, 적절한 δ^2 의 범위는 1에서 10사이이고, C 의 값으로 적합한 범위는 10에서 100사이라고 한다. 이를 참고하여 본 연구에서도 C 와 파라미터의 값을 적합한 범위 내에

서 세분화하여 실험하고, 실험의 결과가 의미 있는 것을 위주로 정리한다. ϵ 은 0.001로 고정한다. 이렇게 실험한 결과, 각 데이터셋마다 약간의 차이가 있지만 커널함수를 가우시안 RBF나 선형으로 사용한 경우에 가장 우수한 성과를 보이고 있음을 알 수 있었다. 아울러 가우시안 RBF일 때 가장 우수한 결과를 보일 때에는, 그 옵션값으로 δ^2 가 1일 때, 가장 우수한 예측정확성을 나타내고 있음도 함께 확인해 볼 수 있었다. 그리하여, 다분류 SVM 알고리즘은 본 채권등급 예측 실험과 관련하여 검증용 데이터에 대해 약 65.50%~68.99% 정도의 예측력을 보이고 있음을 확인하였다.

인공신경망의 경우에는 각 데이터셋마다, 은닉층의 수를 어떻게 설정하는가에 따라 실험결과가 다르게 나타난다. 때문에, 은닉층의 수를 4가지 경우로 변화해 가며 실험을 진행하였으며, 그 결과 인공신경망 모형의 검증데이터에 대한 예측정확성은 각 상품군에 대해 약 64.34%~67.44%인 것으로 나타났다.

〈표 3〉 전체 실험 결과 비교

	MDA		ANN			최적 파라미터 ¹⁾	다분류 SVM		최적 파라미터 ²⁾
	Training	Valid.	Training	Test	Valid.		Training	Valid.	
데이터셋#1	65.86%	59.30%	70.22%	64.73%	64.73%	$h=14$	67.12%	65.89%	MC / LINEAR, $C=10$
데이터셋#2	64.90%	62.02%	71.76%	65.50%	65.12%	$h=7$	84.09%	65.50%	BC / RBF, $C=78$, $\delta^2=1$
데이터셋#3	64.51%	69.38%	69.45%	63.18%	66.67%	$h=21$	67.60%	69.38%	BC / LINEAR, $C=100$
데이터셋#4	64.80%	65.50%	69.19%	64.73%	67.44%	$h=28$	81.20%	68.99%	MC / RBF, $C=55$, $\delta^2=1$
데이터셋#5	65.09%	59.30%	67.14%	68.60%	64.34%	$h=21$	75.31%	68.22%	BC / RBF, $C=10$, $\delta^2=1$
전체 평균	65.03%	63.10%	69.55%	65.35%	65.66%		75.06%	67.60%	

1) h : 은닉층에 속한 노드의 수

2) BC: bound-constrained multi-class support vector classification

MC: multi-class SVC from solving a bound-constrained problem

LINEAR: 선형 (linear) 커널 함수

RBF: 가우시안 RBF (Gaussian RBF) 커널 함수

다음의 <표 3>은 모든 모형의 학습결과를 종합하여 나타난 것이다. 기법 별로 실험을 진행한 결과를 예측정확성을 기준으로 정리하였는데, 인공신경망과 SVM은 여러 가지 파라미터를 조정하여 실험한 결과 중 가장 우수한 결과를 비교하였다.

<표 3>의 결과에서 나타난 것과 같이 예측력은 SVM이 가장 높았고 그 다음으로 인공신경망(ANN), 다변량판별분석(MDA) 순이었다. 특히, 본 실험결과로부터 모든 데이터셋에 대해 항상 SVM의 결과가 우수하다는 것을 알 수 있는데, 이는 기존에 채권등급평가에 적용되어 온 타 알고리즘에 비해 SVM이 얼마나 우수한 지를 입증하는 또 다른 주요한 결과라고 할 수 있다. 평균값으로 볼 때, SVM의 예측력은 다변량판별분석이나 인공신경망 보다 1.94~4.5% 정도 우수한 것으로 나타나고 있는데, 이러한 예측력 차이의 유의성을 검증하기 위하여 McNemar Test를 실시하였다. 그 결과는 <표 4>와 같다.

<표 4> McNemar Test 결과

	ANN	다분류 SVM
MDA	5.044**	12.893***
ANN		3.016*

* 10% 유의수준, ** 5% 유의수준, *** 1% 유의수준

<표 4>에서 나타난 것과 같이 다분류 SVM은 다변량판별분석이나 인공신경망보다 통계적으로 유의한 성과차이를 보이고 있음을 알 수 있다. 아울러 인공지능기법인 인공신경망도 다변량판별분석보다는 5% 유의수준 하에서 통계적으로 유의한 성과차이를 보이고 있음을 알 수 있다. 이를 통해 채권등급평가에 있어서 SVM이 기존의 기법들 보다 우수한 예측정확성을 나타냄을 알 수 있다.

V. 결 론

본 연구에서는 최근 패턴인식 및 분류문제와 관련하여 활발하게 연구되고 있는 SVM을 국내 기업의 채권등급평가에 적용하여 보았다. SVM은 전술한 바와 같이 통계적 이론에 기반하여 설명력이 우수하고, 구조적위험최소화접근에 따라 과도적합문제에서 상대적으로 자유로우며, 불룩함수를 최소화하는 학습을 진행하기 때문에 유일한 최적 또는 최적에 가까운 해를 구할 수 있다는 점이 장점이다. 특히, 본 연구에서는 대표적인 다분류 문제인 채권등급평가분야에 있어서 다변량판별분석, 인공신경망과 비교하여 SVM의 적용가능성을 확인하고자 하였다. 실험 결과, SVM은 상기 기법들보다 우수한 예측력을 보였으며, 다변량판별분석 및 인공신경망과의 예측력 차이가 통계적으로도 유의한 것으로 나타났다. 이와 같이 SVM은 적은 수의 데이터로도 우수한 설명력을 보이는 구조적위험최소화에 기반한 모형을 우리에게 제공해 주는 동시에, 인공신경망과 비슷한 수준의 높은 예측력을 보임으로서, 향후 경영학 분야의 다양한 분류문제에 있어서 유용하게 활용될 수 있을 것으로 생각된다.

특히 본 연구의 대상이 된 채권등급평가 문제에 있어서 SVM의 적용가능성을 검증한 본 연구는 다음의 2가지 측면에서 그 의의를 찾을 수 있다. 우선 첫번째는 본 연구가 부도예측과 같은 이분류 문제뿐만 아니라, 채권등급평가와 같이 다분류 문제에서도 SVM이 매우 우수한 성과를 낼 수 있다는 것을 우리에게 보여주고 있다는 점이다. 전통적인 SVM은 그 기본 원리가 이분류에 기반을 두고 있는 알고리즘으로서, 경영학 분야에서도 지금까지는 주로 부도예측이나 구매예측 등 0 혹은 1의 이분류 문제에 주

로 적용되어 왔다. 하지만, 본 연구는 이러한 이분류 SVM을 확장한 다분류 SVM 방법론도 다분류 문제 해결에 있어서, 이분류 SVM이 타 알고리즘에 비해 보여온 장점을 그대로 보여주고 있으며, 그 결과 가치가 매우 높은 방법론이라는 사실을 보여주고 있다. 때문에 향후 채권분류 뿐만 아니라, 고객 등급 예측과 같은 다른 다분류 경영학 문제의 해결에도 SVM이 좋은 대안이 될 수 있다는 가능성을 본 연구에서 보여주고 있다고 할 수 있다.

둘째로, 본 연구가 채권평가모형 구축의 근본적인 문제인 '자료수의 회소' 문제를 해결할 수 있는 대안으로 다분류 SVM을 제안하고 있다는 점을 들 수 있다. 기존의 분류 문제들과 달리, 채권평가 문제는 다분류 문제로서, 등급별 자료수가 상대적으로 회소하다는 특징을 갖고 있다. 현재까지 인공지능에 기반한 채권분류문제에는 인공신경망이 가장 많이 적용되어 왔는데, 인공신경망은 그 원리상 학습에 많은 데이터를 필요로 하며, 자료수가 충분치 않을 경우 과적합화의 위험이 항상 따르게 된다. 하지만, Shin et al.(2005)의 연구에서 제시된 바와 같이 SVM은 분류기를 도출할 때 경계에 위치한 서포트 벡터들만 선별해 참조하므로, 인공신경망과 달리 과적합화가 거의 발생하지 않는다는 장점이 있다. 따라서, 여러 경영예측 문제 중에서 채권등급평가 분야는 특히 SVM의 적용가능성이 높은 분야임을 본 연구는 시사하고 있다.

본 연구는 이상과 같은 의미를 가지고 있으나 몇 가지 한계점도 가지고 있다. 우선, SVM의 경우, 파라미터를 어떻게 설정하느냐에 따라 성과가 변동될 수 있는데, 아직까지 SVM의 최적 파라미터를 결정할 수 있는 일반적인 방법이 정해져 있지 않는 상황이다. 때문에 본 연구에서는 한정된 범위 내에서만 파라미터 변동의 효과를 확인하고 있는데, 이는

SVM의 우수한 성과를 제대로 제시하지 못하였을 가능성이 있다. 따라서, 이러한 SVM의 파라미터들을 보다 효과적으로 최적화 할 수 있는 연구가 추후 진행되어야 할 필요가 있다.

둘째로, 다분류 SVM의 방법론을 보다 다양화하고, 개선하는 연구도 추가적으로 필요하다. 본 연구에서는 이분류 SVM을 다분류 SVM으로 확장하는 방법론으로 one-against-one 방법론과 Weston & Watkins(1999)의 방법론, 그리고 Crammer & Singer(2000)의 방법론을 적용하였는데, 다분류 SVM 방법론은 이 외에도 다양한 방법이 존재할 수 있다. 특히 채권의 등급평가와 같은 문제의 경우, 다른 다분류 문제와 달리, 등급간 순서가 정해져 있는 특징이 있으므로, 이를 활용하면 보다 효율적이면서도 효과적인 분류가 가능할 수 있다 (Kwon et al., 1997). 앞으로 이러한 다분류 SVM 방법론의 개선에 대해서도 연구자들의 관심과 연구가 필요하다.

끝으로, 다분류 SVM의 다분류 모형에서의 적용가능성을 좀 더 일반화하려면 보다 많은 문제에 적용해 보고, 과연 다른 분야에서도 다분류 SVM이 우수한 성과를 보이는지에 대해 확인해 볼 필요가 있다. 이를 위해서, 다분류 SVM의 다른 경영학 문제에 대한 적용과 검증이 향후 지속적으로 이루어져야 할 것이다.

참고문헌

- 박기남, 이훈영, 박상국 (2000), "러프집합을 이용한 통합형 채권등급 평가모형 구축에 관한 연구," **한국경영과학회지**, 25(3), 125-135.
- 이전창, 한인구, 김명종 (1996), "통계적 모형과 인공지능

- 모형을 결합한 기업신용평가 모형에 관한 연구," *한국경영과학회지*, 21(1), 81-100.
- 이수용, 이일병 (2002), "Fuzzy 이론과 SVM을 이용한 KOSPI 200 지수 패턴분류기," *한국증권학회 제4차 정기학술발표회 논문집*, 787-809.
- 한인구, 권영식, 이진창 (1995), "지능형 기업신용평가시스템의 개발: NICE-AI," *경영학연구*, 24(4), 91-117.
- Altman, E.I., Marco, G., and Varetto, F. (1994), "Corporate Distress Diagnosis Comparisons Using Linear Discriminant Analysis and Neural Networks," *Journal of Banking and Finance*, 18(3), 505-529.
- Chaveesuk, R., Srivaree-Ratana, C., and Smith, A.E. (1993), "Alternative neural network approaches to corporate bond rating," *Journal of Engineering Valuation and Cost Analysis*, 2(2), 117-131.
- Crammer, K., and Singer, Y. (2000), "On the Learnability and Design of Output Codes for Multiclass Problems," *Proceedings of the Thirteenth Annual Conference on Computational Learning Theory*, 35-46.
- Dutta, S. and Shekhar, S. (1988), "Bond rating: a non-conservative application of neural networks," *Proceedings of IEEE International Conference on Neural Networks*, II443-II450.
- Ederington, H.L. (1985), "Classification models and bond ratings," *Financial Review*, 20(4), 237-262.
- Fisher, L. (1959), "Determinants of risk premiums on corporate bonds," *Journal of Political Economy*, 67, 217-237.
- Foody, G.M. and Mathur, A. (2004), "A Relative Evaluation of Multiclass Image Classification by Support Vector Machines," *IEEE Transactions on Geoscience and Remote sensing*, 42(6), 1335-1343.
- Gentry, J.A., Whitford, D.T., and Newbold, P. (1988), "Predicting industrial bond ratings with a probit model and funds flow components," *Financial Review*, 23(3), 269-286.
- Hearst, M.A., Dumais, S.T., Osman, E., Platt, J., and Scholkopf, B. (1998), "Support vector machines," *IEEE Intelligent System*, 13(4), 18-28.
- Hornik, K. (1991), "Approximation capabilities of multilayer feedforward networks," *Neural Networks*, 4, 251-257.
- Hsu, C.W., and Lin, C.J. (2002), "A Comparison of Methods for Multiclass Support Vector Machines," *IEEE Transactions on Neural Networks*, 13(2), 415-425.
- Huang, C.C., and Tseng, T.L. (2004), "Rough set approach to case-based reasoning application," *Expert Systems with Applications*, 26(3), 369-385.
- Huang, Z., Chen, H., Hsu, C-J., Chen, W-H., and Wu, S. (2004), "Credit rating analysis with support vector machines and neural networks: a market comparative study," *Decision Support Systems*, 37, 543-558.
- Jackson, J.D., and Boyd, J.W. (1988), "A statistical approach to modeling the behavior of bond raters," *The Journal of Behavioral Economics*, 17(3), 173-193.
- Jo, H., and Han, I. (1996), "Integration of case-based forecasting, neural network, and discriminant analysis for bankruptcy prediction," *Expert Systems with Applications*, 11, 415-422.
- Joachims, T. (1998), "Text categorization with

- support vector machines," *Proceedings of the European Conference on Machine Learning (ECML), 10th European Conference on Machine Learning*, 137-142.
- Kim, J.W. (1993), "Expert systems for bond rating: a comparative analysis of statistical, rule-based and neural network systems," *Expert Systems*, 10, 167-171.
- Kim, K. (2003), "Financial time series forecasting using support vector machines," *Neurocomputing*, 55(1-2), 307-319.
- Kwon, Y.S., Han, I.G., and Lee, K.C. (1997), "Ordinal Pairwise Partitioning (OPP) approach to neural networks training in bond rating," *Intelligent Systems in Accounting, Finance and Management*, 6, 23-40.
- Lee, K., Han, I., and Kwon, Y. (1996), "Hybrid neural networks for bankruptcy predictions," *Decision Support Systems*, 18, 63-72.
- Maher, J.J., and Sen, T.K. (1997), "Predicting bond ratings using neural networks: a comparison with logistic regression," *Intelligent Systems in Accounting, Finance and Management*, 6, 59-72.
- Osuna, E., Freund, R., and Girosi, F. (1997), "Training support vector machines: an application to face detection," *Proceedings of Computer Vision and Pattern Recognition*, 130-136.
- Patuwo, E., Hu, M.H., and Hung, M.S. (1993), "Two-group classification using neural networks," *Decision Science*, 24(4), 825-845.
- Pinches, G.E., and Mingo, K.A. (1973), "A multivariate analysis of industrial bond ratings," *Journal of Finance*, 28(1), 1-18.
- Pinches, G.E., and Mingo, K.A. (1975), "The role of subordination and industrial bond ratings," *Journal of Finance*, 30(1), 201-206.
- Pogue, T.F., and Soldofsky, R.M. (1969), "What's in a bond rating?," *Journal of Financial and Quantitative Analysis*, 4(2), 201-228.
- Ramaswamy, S., Tamayo, P., Rifkin, R., Mukherjee, S., Yeang, C-H., Angelo, M., Ladd, C., Reich, M., Latulippe, E., Mesirov, J.P., Poggio, T., Gerald, W., Loda, M., Lander, E.S., and Golub, T.R. (2001), "Multiclass cancer diagnosis using tumor gene expression signatures," *Proceedings of PNAS* 98, 15149-15154.
- Shin, K.-S., and Han, I. (1999), "Case-based reasoning supported by genetic algorithms for corporate bond rating," *Expert Systems with Applications*, 16(2), 85-95.
- Shin, K.-S., and Han, I. (2001), "A case-based approach using inductive indexing for corporate bond rating," *Decision Support Systems*, 32, 41~52.
- Shin, K.-S., Lee, T.S., and Kim, H.-j. (2005), "An application of support vector machines in bankruptcy prediction model," *Expert Systems with Applications*, 28(1), 127-135.
- Statnikov, A., Aliferis, C.F., Tsamardinos, I., Hardin, D., and Levy, S. (2005), "A comprehensive evaluation of multicategory classification methods for microarray gene expression cancer diagnosis," *Bioinformatics*, 21, 631-643.
- Tay, F.E.H., and Cao, L.J. (2001), "Application of support vector machines in financial time series forecasting," *Omega*, 29, 309-317.
- Tay, F.E.H., and Cao, L.J. (2002), "Modified support vector machines in financial time

- series forecasting," *Neurocomputing*, 48, 847-861.
- Vapnik, V. (1995), *The Nature of Statistical Learning Theory*, Springer-Verlag, New York.
- West, R.R. (1970), "An alternative approach to predicting corporate bond ratings," *Journal of Accounting Research*, 8(1), 118-125.
- Weston, J., and Watkins, C. (1999), "Support Vector Machines for Multiclass Pattern Recognition," *Proceedings of the Seventh European Symposium on Artificial Neural Networks*, 219-224.

Intelligent Credit Rating Model for Korean Companies using Multiclass Support Vector Machines

Hyunchul Ahn* · Kyoung-jae Kim** · Ingoo Han***

Abstract

Investors, debt issuers, and others are quite interested in corporate credit rating (i.e. bond rating) because it is considered to be an important measure for managing financial risk of their portfolios. But, in spite of its importance, corporate credit rating is typically very costly to obtain, since it requires professional agencies to invest large amount of time and human resources to perform deep analysis of the company's risk status based on various aspects ranging from strategic competitiveness to operational level details. As a result, it has been a popular research topic for researchers to predict companies' credit ratings by applying statistical and artificial intelligence techniques

The researchers in the early days mainly focused on applicability of statistical techniques such as multiple discriminant analysis (MDA) and logistic regression (LOGIT) analysis. However, more recent studies have shown that artificial neural networks (ANNs) achieved better performance than traditional statistical methods and other artificial intelligence methods in bond rating. Consequently, ANN has been the most widely-used technique for corporate credit rating for a long time.

However, despite ANN's superior performance, it has some critical limitations. First of all, it suffers from difficulty in selecting a large number of controlling parameters which include relevant input variables, hidden layer size, learning rate, and momentum term. And, it generally requires a large data set for the effective training. Furthermore, there is a danger of overfitting, and ANN usually requires huge amount of data samples for the effective

* Postdoc Researcher, Graduate School of Management, Korea Advanced Institute of Science and Technology

** Assistant Professor, Department of Management Information Systems, Dongguk University

*** Professor, Graduate School of Management, Korea Advanced Institute of Science and Technology

training. In addition, it is usually difficult to explain why it produces a specific result, i.e. poor explanation ability.

To mitigate these limitations, this paper suggests a novel machine learning technique, multi-class support vector machine (SVM), as a tool for corporate credit rating. General SVM is simple enough to be analyzed mathematically, and leads to high performance in practical applications. Although many traditional neural network models have implemented the empirical risk minimization principle, SVM implements the structural risk minimization principle. The former seeks to minimize the misclassification error or deviation from the correct solution of the training data but the latter searches to minimize an upper bound of the generalization error. In addition, the solution of SVM may be a global optimum, while other neural network models may tend to fall into a local optimal solution. Thus, overfitting is unlikely to occur with SVM. In addition, SVM does not require too many data samples for training since it builds prediction models by only using some representative samples near the boundaries (so-called *support vectors*).

However, original SVMs were originally devised for binary classification. Thus, in order to apply them to multi-class classification problems such as corporate credit rating, the original SVMs should be extended to the multi-class SVM models. So far, researchers have proposed various approaches for this extension. In this study, we have tried three different methods for multi-class SVMs including (1) bound-constrained multi-class support vector classification (SVC), (2) multi-class SVC from solving a bound-constrained problem, and (3) multi-class SVC from Crammer and Singer. And, we try to find the most appropriate method and parameters for our data set.

To examine the feasibility of multi-class SVMs in corporate credit rating, we applied these methods to the real-world bond rating case for Korean companies. We applied multi-class SVMs as well as ANN and MDA to the same dataset, so we tried to validate the superiority of multi-class SVM to other comparative algorithms. The experimental results showed that multi-class SVM outperformed ANN and MDA with a statistical significance level of 0.05 to 0.10. As a result, we found that multi-class SVM is a promising alternative to corporate credit rating prediction.

Key words: Multi-class support vector machine, Bond rating, Artificial neural networks, Multivariate discriminant analysis