

RDB データベースにおける信用リスクモデルの 説明力の年度間推移に関する分析

柳澤 健太郎 下田 啓 岡田 絵理 清水 信宏 * 野口 雅之 †

概 要

金融機関において信用リスクモデルの構築のみならず、モデルの検証に対して関心が高まっていることの一つの理由として、2007 年 3 月末から実施された新しい自己資本規制（バーゼル II）がある。バーゼル II では内部格付制度におけるモデルの利用に関して、利用するモデルの予測能力（説明力）が高いこと、モデルの運用実績及び安定性の評価を定期的に行うことが要求される。モデル構築時に比べ検証時の説明力が推移している場合はその要因を究明する必要がある。

本分析では RDB(Risk Data Bank) データベース¹を使用して複数の信用リスクモデルを構築し、モデルの説明力の推移について検証を行った。その結果、いずれのモデルにおいても 2000 年度～2006 年度の RDB データについて説明力の傾向的な低下が確認された為、その要因の究明を行い、説明力の低下がデフォルトスコアの高スコア側へのシフトに起因する事を示した。以上の分析は、モデル性能の安定性確保のための対応を考える際にも有効と考えられる。

1 はじめに

金融機関では債務者評価の 1 つの手段として信用リスクモデル²が広く利用されているが、バーゼル II の導入によりモデル性能の評価手法について関心が高まっている。バーゼル II では内部格付制度でモデルを利用する場合には、モデルの予測能力（説明力）が高く、モデルの検証を定期的に行うこと等が求められる。

モデルの説明力を測る手段としてはさまざまな統計量が存在するが、代表的なものは AR(Accuracy Ratio) 値 [1][13] であり、監督当局はこの値を重視していると言われている。AR 値に関しては、過去に行ったモデル構築とその検証等を通じてその水準が大きく変化する場合があることがわかっている。AR 値に関する研究として Hamerle, Rauhmeir, Rosch(2003)[1] があり、その中では AR 値のような統計量は、モデルの説明力だけではなく計測対象としているポートフォリオの構造の影響を受けるので、値の解釈には注意が必要であるとされている。

本分析では、時点の異なるデータを利用して複数のモデルを構築し、各年度のデータで検証を行うことにより、AR 値の変動要因を確認した。手順としては、最初に 2000 年度から 2006 年度までの年度ごとのデータを用い、計 7 本のモデルを作成した。次に作成したモデルごとに、各年度

*日本リスク・データ・バンク株式会社

†データ・フォアビジョン株式会社

¹3 大金融グループおよび地方銀行を中心とした金融機関約 50 行によってプールされているデータ。金融機関の融資先である非上場企業（中堅・中小企業）に関するステータス情報や財務情報、顧客情報（企業名は除く）が含まれる。

²ここでは債務者の信用リスク（デフォルト確率等）を評価するモデルを指す。

のデータに対して AR 値を測定し、水準を比較した。

この結果、いずれのモデルにおいても AR 値が経年変化していることが確認されたため、その要因がモデルの説明力の低下によるものか、ポートフォリオの構造変化によるものかを分析した。AR 値の変動が特定のモデルに限った現象であれば、要因はそのモデルの説明力低下にあり、どのモデルに対しても共通に観測される現象であれば、要因はポートフォリオの構造変化にあると判断されたと考えた。

今回の分析には RDB データを利用した。分析対象は 2000 年度から 2006 年度の間に非デフォルトまたはデフォルトのいずれかに認定された債務者のうち、ステータス認定の直前期に財務が存在する債務者とした。デフォルトの定義は 3ヶ月以上の延滞または債務者区分が破綻懸念先以下に認定された場合とした。

本稿の構成は以下のとおりである。2 節では本分析で用いたモデルの特徴と、モデルの説明力を表す指標として用いた AR 値の意味を解説し、3 節では検証方法、4 節ではモデルの構築手順を説明する。5 節で分析結果を示す。そして、最後に 6 節は分析結果をまとめる。

2 信用リスクモデルと説明力

ここでは、信用リスクモデルにロジスティック回帰モデル [9] を使用し、その説明力を表す統計量として AR 値 [1] を用いる。信用リスクモデルの構築手法と説明力を表す統計量は多数存在し、それぞれに優劣つけがたく、どの手法を用いるかは議論の分かれるところではあるが、一般的に使われているロジスティック回帰モデルと AR 値をここでは採用した。各々の詳細については以下の小節に記す。

2.1 ロジスティック回帰モデル

本論で用いる信用リスクモデルでは、企業の財務指標からデフォルト率を推定する。そのモデルの構築手法として、ロジスティック回帰モデルをここでは使用する。観測された変数群として $\mathbf{x} = (x_1, \dots, x_r)$ が与えられた時、ロジスティック回帰モデルによりこの条件の下でデフォルトが発生しない条件付確率 $p(\mathbf{x})$ は以下の式で与えられる。

$$p(\mathbf{x}) = \frac{\exp(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_r x_r)}{1 + \exp(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_r x_r)} \quad (1)$$

ここで $\beta = (\beta_0, \dots, \beta_r)$ はパラメータ群である。一般的には、 $p(\mathbf{x})$ をデフォルトが発生する条件付確率とするが、ここではスコア $\beta(\beta = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_r x_r)$ が高い程デフォルト率が低下する様に設定する為、デフォルトが発生しない条件付確率と定義する。また、本論中の議論では主としてデフォルト率ではなくスコアを使用する。

パラメータ群 $\beta = (\beta_0, \dots, \beta_r)$ は最尤法により求める。尤度関数は以下で与えられる。

$$like(\beta) = \prod_{i=1}^I p(\mathbf{x}_i)^{y_i} (1 - p(\mathbf{x}_i))^{1-y_i} \quad (2)$$

ここで、 I はパラメータ推定に用いる顧客数を、 y_i は i 番目の顧客のデフォルト状態 ($y_i = 0$; デフォルト, $y_i = 1$; 非デフォルト) を、 x_i は i 番目の顧客の観測された変数群を表す。この尤度関数を最大とする β を見つけることで、パラメータ群の値が決定する。なお、ロジスティック回帰モデルのパラメータ推定には SAS/STAT® の LOGISTIC プロシジャを使用した。

2.2 AR 値

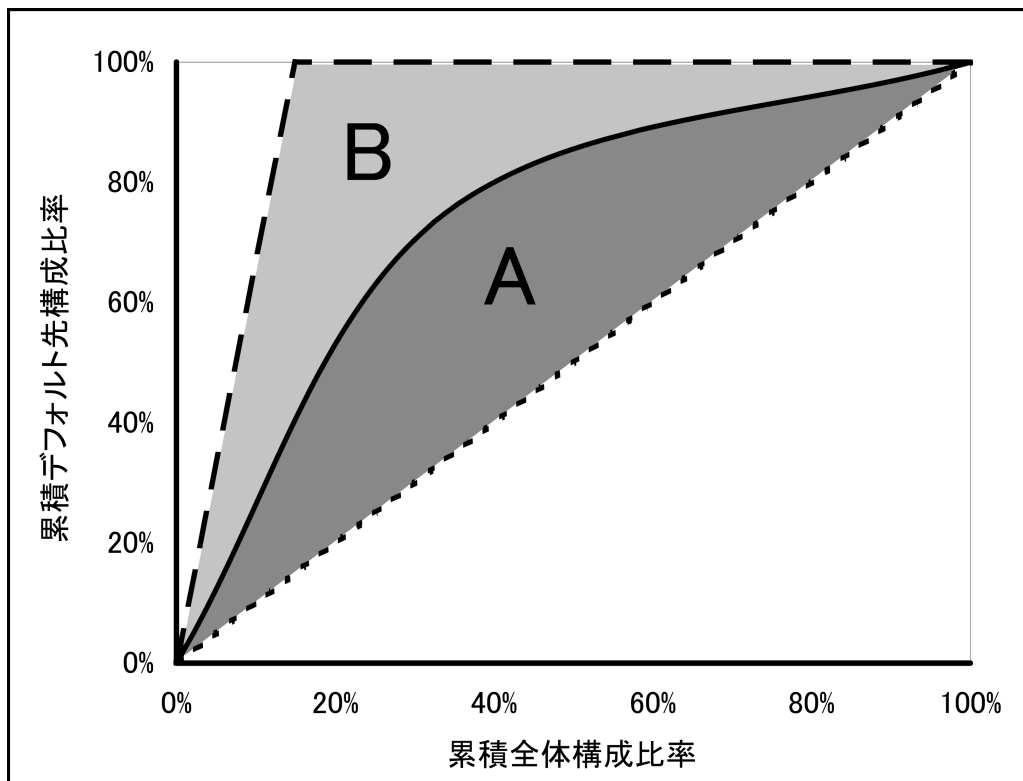


図 1: CAP 曲線と AR 値

信用リスクモデルの説明力を表す統計指標にここでは AR 値を用いる。AR 値の説明の為に、まず CAP 曲線 (Cumulative Accuracy Profiles) について記す。ロジスティック回帰モデルによって推定されたデフォルト率の高い順、つまりスコアの低い順に債務者を並べ、カットオフ値をスコアの小さい値から大きい値に順次変化させ、カットオフ値以下の債務者数の全債務者数に対する割合を x 軸に、カットオフ値以下のデフォルトの債務者数の全デフォルト債務者数に対する割合を y 軸にとったものが CAP 曲線であり、図 1 の実線で表される様な曲線となる。債務者につけられたスコアが完全に正しい場合は、スコアの低い債務者から順次デフォルトする為、図 1 の破線の CAP 曲線 (実線の上側) で表される。また、債務者につけられたスコアの意味が完全でない場合は、図 1 の点線の CAP 曲線 (実線の下側) で表される。図 1 の網掛けされた箇所の面積を A、B とおく場合、AR 値は以下の式で与えられる。

$$AR = \frac{A}{A + B} \quad (3)$$

式からも分かるように、AR 値は推定されたスコアの序列の、完全に正しくつけられたスコアの序

列からの乖離率を表している。

なお、本論中では CAP 曲線下の面積は台形法を用いて算出した。

3 検証手順

本分析では、RDB データベースに含まれる 2000 年度～2006 年度まで計 7 年間のデータを使用し、各年度のデータだけを用いて 7 種類のモデルを作成する。その作成した 7 種類のモデルごとに、2000 年度～2006 年度の各年度のデータに対して計 49 個の AR 値を測定し、その水準を比較する。

本論では、説明力の年度間推移について分析を行いたい訳であるが、異なる年度のデータを用い構築した複数のモデルの AR 値を比較する事で、特定のモデルが特定の年度のデータに著しくオーバーフィットしていないかを確認できる。また、副題として信用リスクモデルの性能が、モデル構築に使用したデータ年度に依存するかを調査する。その為、本分析では異なる年度のデータを用い 7 種類のモデルを作成した。

4 モデル構築手順

本節では、ロジスティック回帰モデルの本論での構築手順について記す。本分析で構築した 7 種類のモデルは、モデル毎に最尤法による係数の最適化だけでなく、AR 値が最大になるように説明変数を最適化した。また、モデル間の比較を行う際にモデル構築によるバイアスを最小限に抑える為に、同じモデル構築条件の下、以下のロジックに従い機械的にモデルを構築した。

1. 説明変数の候補となる財務指標の選択
2. Box-Cox 変換による財務指標値の変換
3. SGA による説明変数の選択
4. 最尤法によるパラメータ群の推定

以下の小節にその詳細を記す。

4.1 候補となる財務指標の選択

RDB データベースには 106 の財務項目が定義されている。本分析では、その 106 の財務項目から一般的に使われている 91 の財務指標を作成し、ロジスティック回帰モデルの説明変数の候補とした。

候補指標の中には、分布が大きく裾を引いているものがある。例えば、2003 年度の RDB データベースのクレジット・インターバルの記述統計量は、平均値 = 1.90, 標準偏差 = 7.45, 最小値 = 0, 5%点 = 0.10, 中央値 = 1.02, 95%点 = 5.50, 99%点 = 13.47, 最大値 = 1157.54 である。この事から、2003 年度のクレジット・インターバルの分布は大きく右に裾を引いている事が分かる。また、最大値は明らかに外れ値である。これらの外れ値に近い値の分析への影響を緩和

する為に、候補指標値に上限値(下限値)を設け、その上限値を超える(下回る)値は上限値(下限値)に置き換えている。なお、下限値には5%点の値を、上限値には95%点の値を使用した。

候補指標の中には上下限値や0の値を持つ債務者が多い財務指標も存在する。そのような財務指標がロジスティック回帰モデルの説明変数に選ばれた場合、債務者を正しく評価できない可能性が高くなる。そこで、上下限値、もしくは0の値をもつ債務者数が全債務者の30%を超える財務指標に関しては、候補指標から除外した。その結果、候補指標から9つの財務指標が除外され82となった。82の候補指標の一覧を表1に記す。

表 1: 候補指標一覧

キャッシュレシオ	インターバルメジャー	借入金自己資本比率	実質借入金回転期間
実質有利子負債回転期間	支払利息割引料総利益率	手元流動性比率	売上高金利負担率
売上高純金利負担率	金利カバレッジ	支払利息割引料現金預金率	現預金比率
売上高キャッシュフロー率	総資本キャッシュフロー率	総負債キャッシュフロー率	有利子負債キャッシュフロー率
総キャッシュフロー比率	クレジット・インターバル	総支出キャッシュフロー率	経費キャッシュフロー率
売上高付加価値率	労働分配率	設備投資効率Ⅱ	減価償却費率
損益分岐点比率	経常損益比率	総資産現金預金率	総資産
自己資本	運転資本比率	当座比率	売上高運転資本比率
買入債務売上債権率	自己資本比率	固定比率	固定長期適合率
借入金依存度	長期負債比率	負債比率	総資本資本金率
総負債資本金率	総キャピタリゼーション比率	売上高借入金比率	クイックレシオ
純資産倍率	総資産事業資本率	売上高総利益率	売上高販管費率
売上高営業利益率	売上高経常利益率	売上高税引前当期利益率	売上高当期利益率
総資本総利益率	総資本営業利益率	総資本経常利益率	総資本当期利益率
ROA	総資本当期末処分利益率	流動比率	自己資本総利益率
自己資本営業利益率	自己資本経常利益率	自己資本当期利益率	現金預金短期借入金率
総資本回転率	固定資産回転率	売上債権回転期間	棚卸資産回転期間
買入債務回転期間	借入金回転期間	立替期間	現金預金有利子負債率
有利子負債構成比率	デット・エクイティ・レシオ	有利子負債EBITDA率	デット・キャパシティ・レシオ
有利子負債回転期間	有利子負債利率	総資産実質借入金率	総資産実質有利子負債率
資本金	売上高		

4.2 財務指標値の変換

財務指標をロジスティック回帰モデルの説明変数にするにあたり、値の変換を行った。変換手法には Box-Cox 変換 [3] と正規分布の累積確率への変換を用いた。

Box-Cox 変換とは観測データ(ここでは財務指標値)を正規分布に近い分布に変換する手法である。Box-Cox 変換は以下の式となる。

$$\begin{aligned}
 z &= \left((y+c)^\lambda - 1 \right) / \lambda & \text{if } \lambda \neq 0 \\
 z &= \ln(y+c) & \text{if } \lambda = 0
 \end{aligned} \tag{4}$$

c は正の実数であり、 y が負値の場合、正值にリスケールするパラメータである。また、 λ は実数であり変換のパラメータである。つまり、変換値の分布が最も正規分布に近くなるように、 λ を決

定する．なお， λ の推定には SAS/STAT[®] の TRANSREG プロシジャを使用した．

次に，財務指標値の Box-Cox 変換後に正規分布の累積確率への変換を行った．これは，全ての財務指標値のスケールを一致させ，値の範囲を 0 以上 1 以下にする為である．以下の正規分布の累積分布関数を用い，正規分布の累積確率への変換を実施した．

$$\int_{-\infty}^z \frac{1}{\sqrt{2\pi}\sigma} \exp \left[-\frac{(x-\mu)^2}{2\sigma^2} \right] dx \quad (5)$$

ここで， μ と σ^2 は，Box-Cox 変換後のデータでの平均と分散を表す．

4.3 説明変数の選択

SGA(Simple Genetic Algorithm: 単純遺伝的アルゴリズム)[8] を使用し，ロジスティック回帰モデルの AR 値が最大になるように候補指標から説明変数の組み合わせを選択した．なお，説明変数は 10 個を目安に選択した．遺伝的アルゴリズムとは最適化問題を解くための手法の一つであり，生物の進化のメカニズムを模擬する人工モデルとして提唱されてきた手法とみなされている．そして，基本的な遺伝的アルゴリズムを SGA と呼ぶ．

説明変数を選択する手法としては，「ステップワイズ法」，「前進法」，「後退法」[9] が一般的である．しかし，「ステップワイズ法」では説明変数の数を任意に指定できない欠点がある．また，これらの選択法では一般的に尤度比検定やスコア検定，Wald 検定などを用いて最適な変数の組み合わせを求める．本分析では説明変数の数を任意に指定し，AR が最大となる説明変数の組み合わせを得る為に，SGA を使用した．一方，SGA にも欠点がある（最適化計算に時間がかかる，一般的なパッケージでは計算できない等）．その為，どの選択法を選ぶかについては問題に応じて考えるべきであろう．

次に，SGA を用いた最適化手順を以下に示す．

1. 個体群中の各個体の適合度を求める
2. 適合度の高い個体が高確率で生き残るように再生する
3. 交叉と突然変異を起こす
4. エリートを個体群に加える

SGA では，手順 1～4 を繰り返し実施し，個体群の世代を重ねることで最適解へと近づいて行く．そして，定義した終了基準を満たした段階で繰り返しを終了する．

本分析では，SGA を説明変数の最適な組み合わせ選択に用いる為に，適合度として AR 値を，個体は説明変数の組み合わせとした．その他の条件は，再生の選択法はルーレット選択，交叉は一点交叉，個体数は 50，世代ギャップ（入れ替える新個体の割合）は 50%，交叉確率は 60%，突然変異確率は 0.5%，終了基準として繰り返し回数を 500 回と設定した．また，説明変数の数が 10 個から外れる場合には，外れた度数に応じて適合度にペナルティを与えた．SGA は，SAS[®] のデータステップ等を使用した．なお一般的には，説明変数の選択時には交互作用項の検討を行うが，本

分析への影響は小さいものと考え、交互作用項の検討を行わなかった。

SGA より得られた説明変数の組み合わせは、7種類のモデル全てで異なった。例えば、2003年度データによって選択された説明変数は、借入金自己資本比率、支払利息割引料総利益率、支払利息割引料現金預金率、クレジット・インターバル、売上高付加価値率、経常損益比率、純資産倍率、総資産事業資本率、総資本回転率、固定資産回転率、有利子負債利子率の11指標であった。総資産事業資本率は7種類全てのモデルの説明変数に選択された。

4.4 パラメータ群の推定

4.3節で得た説明変数の組み合わせをロジスティック回帰モデルに当てはめ、パラメータ群の推定を行った。この時、全ての説明変数についてWald検定のp値が有意水準1%で有意である事を確認した。また一般的には、モデル構築時にパラメータの推定値の方向(正負)が、説明変数の特性と一致しているかを確認するが、本分析ではその手順を省略した(例えば、自己資本比率の場合、大きいほど優良な企業であるのでパラメータの推定値が正となる必要がある)。パラメータの推定にはSAS/STAT®のLOGISTICプロシジャを使用した。

5 説明力の年度間推移

本節では4節で構築したモデルを使用し、RDBデータベースから得られた結果を示す。小節5.1では、7種類のモデルのAR値の年度間推移について、小節5.2では、2003年度のデータを用い作成したモデルのスコア値の年度間推移について分析を行う。他のモデルについてもスコア値の年度間推移の傾向は同様であったため、データの間地点にあたる2003年度のデータを用い作成したモデルの結果のみを示した。

5.1 AR 値の年度間推移

構築した7種類のモデルの7種類のデータに対するAR値は表2の通りであった。

表 2: 7種類のモデルにおけるAR値の年度間推移

計測対象 データ	モデル(作成データ年度)						
	2000	2001	2002	2003	2004	2005	2006
2000年度	71.72%	69.33%	70.32%	70.84%	69.52%	68.78%	69.43%
2001年度	71.70%	73.46%	72.97%	72.61%	71.74%	70.71%	71.98%
2002年度	68.00%	68.07%	69.53%	68.64%	67.32%	66.86%	67.81%
2003年度	66.35%	65.05%	66.00%	67.49%	66.38%	66.22%	65.91%
2004年度	61.10%	60.50%	61.23%	61.87%	62.72%	61.93%	60.22%
2005年度	63.04%	62.07%	62.79%	64.15%	64.43%	65.12%	63.89%
2006年度	59.41%	60.10%	60.45%	61.34%	61.02%	62.07%	62.62%

計測対象データとは AR 値の測定に使用した RDB データベースのデータ年度を指す。表 2 の行は同じ計測対象データに対する各モデルの AR 値に対応している。また、列は同じモデルの異なる計測対象データに対する AR 値を表している。網掛けはインサンプルデータ、すなわちモデル構築に使用したデータに対応する。網掛けしていないセルはアウトサンプルに対する AR 値となる。同じ計測対象データ内で、インサンプルデータの AR 値が高めに計測されているが、これはモデルを構築したデータを用い AR 値を計測している為である。表 2 の AR 値をモデル毎にプロットし、直線で結んだものが図 2 である。

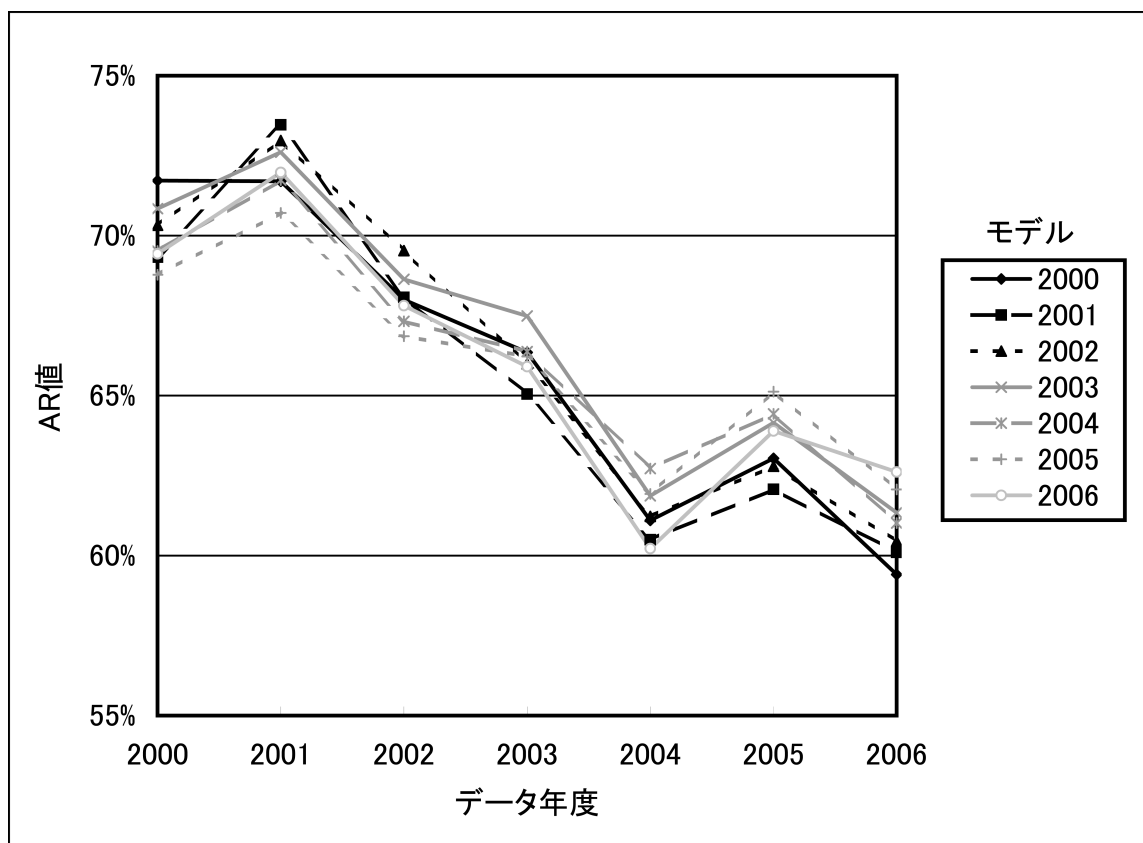


図 2: 7 種類のモデルにおける AR 値の年度間推移

表 2 と図 2 から分かるように、7 種類全てのモデルで、計測対象データに対して時系列に AR 値が低下傾向にあり、2000 年度から 2006 年度の 7 年間で、AR 値が約 10% 低下している。また、期間を通して比較すると、どのモデルも優劣つけ難く、特定のモデルが特定の年度のデータに著しくオーバーフィットしていない事がわかる。

次に、モデル内の AR 値の比較の為に表 3 を、モデル間の AR 値の比較の為に表 4 を作成した。表 3 は同一モデル内 (異なる計測対象データ間) における AR 値の順位を、表 4 は同一計測対象データ内 (異なるモデル間) における AR 値の順位を表している。

表 3: 同一モデル内での AR 値の順位

計測対象 データ	モデル(作成データ年度)						
	2000	2001	2002	2003	2004	2005	2006
2000年度	1	2	2	2	2	2	2
2001年度	2	1	1	1	1	1	1
2002年度	3	3	3	3	3	3	3
2003年度	4	4	4	4	4	4	4
2004年度	6	6	6	6	6	7	7
2005年度	5	5	5	5	5	5	5
2006年度	7	7	7	7	7	6	6

表 3 では、モデル内で一番 AR 値の高い計測対象データに網掛けをした。どのモデルであっても同じ計測対象データに対しては、ほぼ同じ順位である。また、全てのモデルにおいて AR 値の順位は年を追うごとに下がる傾向にある。このことから、モデルの違いに関わらず AR 値は直近に向かって低下する傾向にある事が分かる。

表 4: 同一計測対象データ内での AR 値の順位

計測対象 データ	モデル(作成データ年度)						
	2000	2001	2002	2003	2004	2005	2006
2000年度	1	6	3	2	4	7	5
2001年度	6	1	2	3	5	7	4
2002年度	4	3	1	2	6	7	5
2003年度	3	7	5	1	2	4	6
2004年度	5	6	4	3	1	2	7
2005年度	5	7	6	3	2	1	4
2006年度	7	6	5	3	4	2	1

表 4 では、モデル間で一番 AR 値の高い計測対象データに網掛けをした。その網掛けが対角線上にあり、一番 AR 値の高いモデルとインサンプルデータで作成したモデルが完全に一致している。アウトサンプル(網掛けのしていない)データに対するモデル間の AR 値の比較から、明らかに性能の高いモデルは存在しない。しかし、2003 年度データで構築したモデルの AR 値が全体的に高いようにも見える。また、対角線上の上下の AR 値が高いようにも見える。

ここで、表 4 の特徴から「モデル構築に使用したデータから離れた年度の計測対象データほど、AR 値が低い傾向にある」という仮定を置く。この仮定の下では、2003 年度(本分析に使用した中間地点にあたる)のデータで構築したモデルでは、全体的に計測対象データとの年度差が小さいた

め、AR 値が高い傾向にあると説明できる。また、インサンプルデータの前後で AR 値が高い事の説明もつく。この仮定を検証するために、計測対象データ毎に回帰分析 [11] を実施する。回帰分析の目的変数には AR 値を、説明変数には計測対象データとモデル構築に使用したデータの年度の差の絶対値を使用する。また、回帰分析ではインサンプルデータを除外し解析を実施した。例えば、2003 年度データに回帰分析を実施する場合、目的変数と説明変数の値の組み合わせは表 2 より、(66.35, 3), (65.05, 2), (66.00, 1), (66.38, 1), (66.22, 2), (65.19, 3) の 6 通りとなる。

回帰分析の結果を、表 5 に示す。有意水準を 5% とした場合、傾きの p 値に関しては、2006 年度の計測対象データのみが有意である。つまり、2006 年度の計測対象データに対しては、「モデル構築に使用したデータから離れた年度の計測対象データほど、AR 値が低い傾向にある。」という仮定が成り立っているが、他の計測対象データに対しては成立していない。故に、データ全体で見た場合にはこの仮定が成立しているとは言えない。

表 5: 計測対象データ毎の回帰分析

計測対象 データ	切片				傾き			
	推定値	標準誤差	t統計量	p値	推定値	標準誤差	t統計量	p値
2000年度	70.25	0.71	98.45	0.0000	-0.16	0.18	-0.85	0.4440
2001年度	72.63	0.63	115.33	0.0000	-0.25	0.21	-1.23	0.2865
2002年度	68.42	0.53	128.05	0.0000	-0.29	0.22	-1.33	0.2537
2003年度	66.05	0.60	110.67	0.0000	-0.03	0.28	-0.11	0.9187
2004年度	61.85	0.60	103.07	0.0000	-0.33	0.25	-1.32	0.2572
2005年度	64.57	0.52	123.28	0.0000	-0.44	0.17	-2.56	0.0627
2006年度	62.43	0.28	221.94	0.0000	-0.48	0.07	-6.71	0.0026

以上の結果から、二つの知見を得た。

1. 2000 年度から 2006 年度の RDB データベースの AR 値は時系列に低下傾向にある
2. 2000 年度から 2006 年度のいずれのデータでモデルを構築しても、モデルの性能に大差はない

ただし、この知見にはデータ件数について注意点がある。本分析において、モデル構築、AR 値測定に用いたデータ件数は、年度間で多少の差があるものの、概ねデフォルトが 5,000 件、非デフォルトが 30,000 件程度ある。件数が少なくなることで、モデルはインサンプルデータによりオーバーフィットし、測定される AR 値の分散は大きくなる。その為、件数が少ない場合 (特にデフォルトが 1,000 件に満たない場合) は、これらの誤差の影響を考慮する必要がある。また、期間 (2000～2006) が増加した場合には違った傾向が確認される可能性もある。本分析では、傾きの p 値の有意性が 2006 年度の計測対象データのみにおいて確認された。その有意性が偶然発生したものなのか、計測対象データとモデル構築データの年度が離れた影響なのかは、本分析では明らかにできなかった。その為、今後も継続して分析を実施し、期間が増加した場合の変化について調査する必要がある。

5.2 スコアの年度間推移

本小節では、スコアの年度間推移について分析を行う。なお、スコアの年度間推移についてはデフォルト、非デフォルト毎に分析を実施する。また、スコアの計算には2003年度データで作成したモデルを使用した。

まず始めに、図3に一年置きのスコアの分布の年度間推移を示す。左側の山がデフォルト、右側の山が非デフォルトを表している。スコアの分布はデフォルト、非デフォルトごとに、スコアの0.5ポイント刻みに密度を求め、各点をスプライン曲線で繋いで描いた。図3では、デフォルト側の分布が時系列で右側、つまり高スコア側へ移動している。また、非デフォルトの分布の移動に比べ、デフォルトの分布の移動が大きく、年を追う毎にデフォルトと非デフォルトの分布の重なりが大きくなっている。

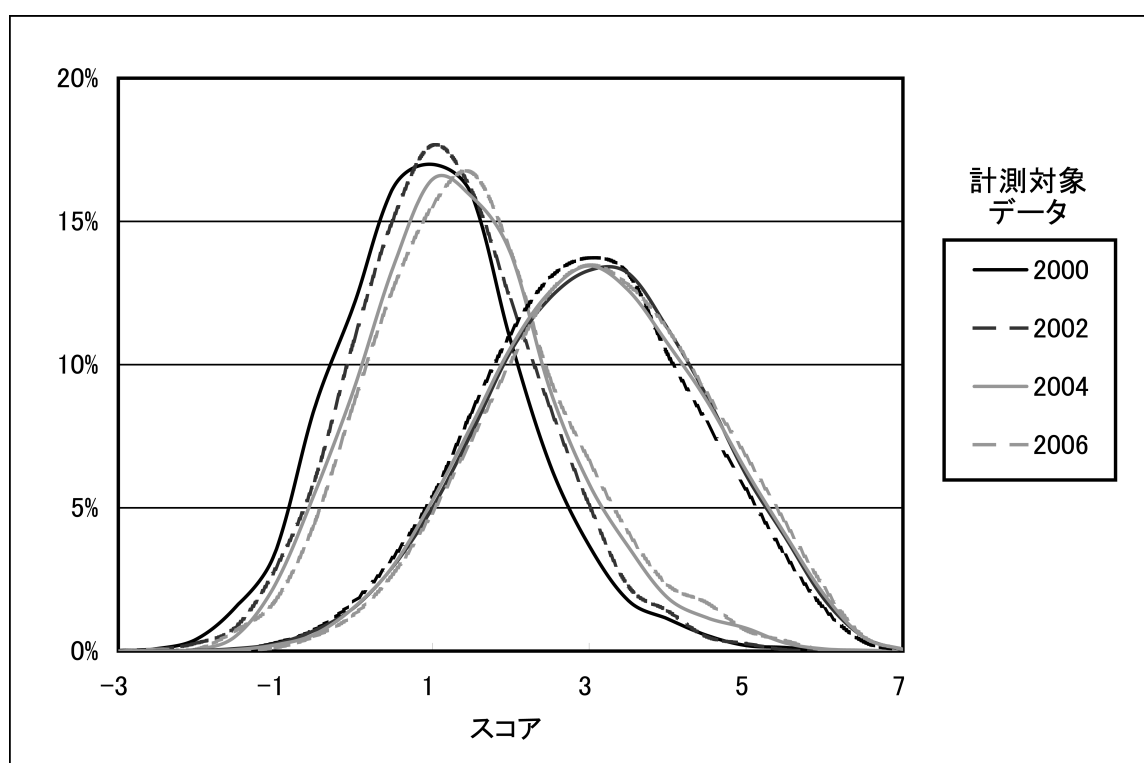


図 3: スコアの分布の推移

次に、スコアの平均値の年度間推移の傾向を把握する為、年度ごとのデフォルト、非デフォルトそれぞれのスコアの平均値について、全体の平均値との乖離をとり、その年度間の推移を図4に示した。デフォルトは傾向を持って高スコア側へ移動しているが、非デフォルトは上昇と下降を繰り返している。つまり、スコアの平均値の推移がデフォルトと非デフォルトでは異なる傾向にある。

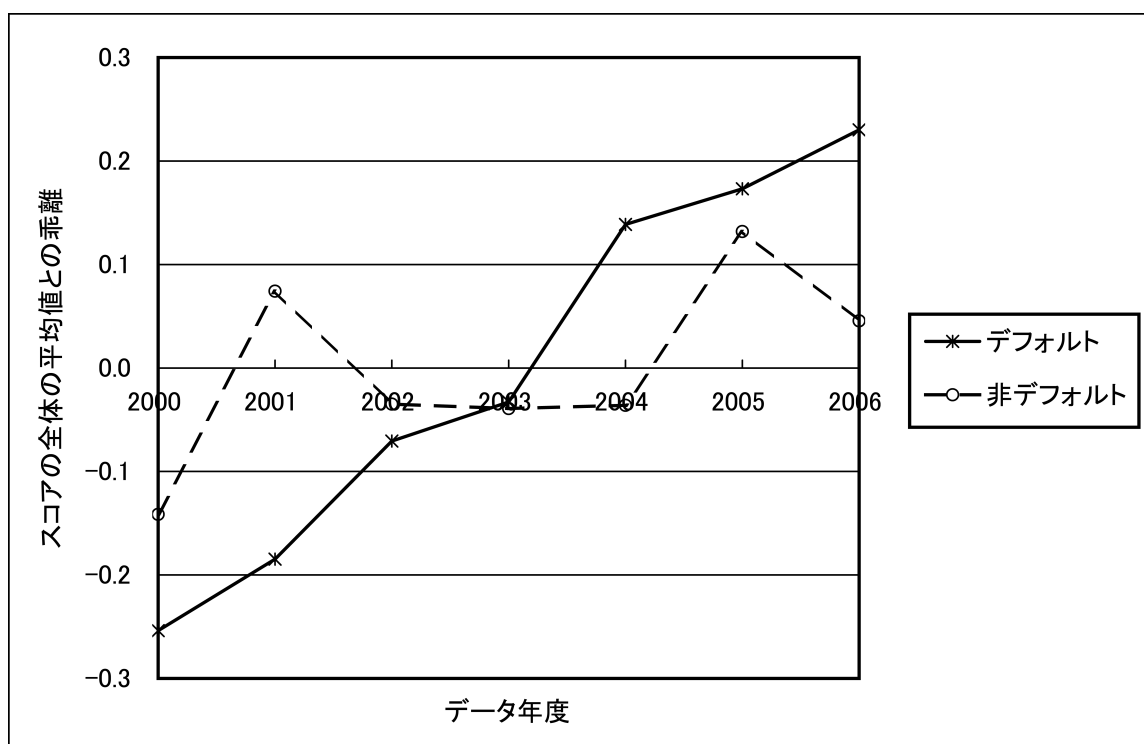


図 4: スコアの全体の平均値との乖離の年度間推移

次に、デフォルト、非デフォルトのスコアの平均値について、データ年度に対する傾きを最小二乗法を用い推定した回帰直線と併せて図5に出力した。回帰直線の傾きの大きさは、デフォルトでは0.0868、非デフォルトでは0.0279であった。デフォルト、非デフォルトともにスコアの平均値が高スコア側に移動する傾向があるが、デフォルトの傾きは非デフォルトに比べて約3倍大きく、デフォルトの分布と非デフォルトの距離が近づくことで、デフォルトの分布と非デフォルトの分布が重なる部分が大きくなった結果、AR 値が低下しているものと考えられる。

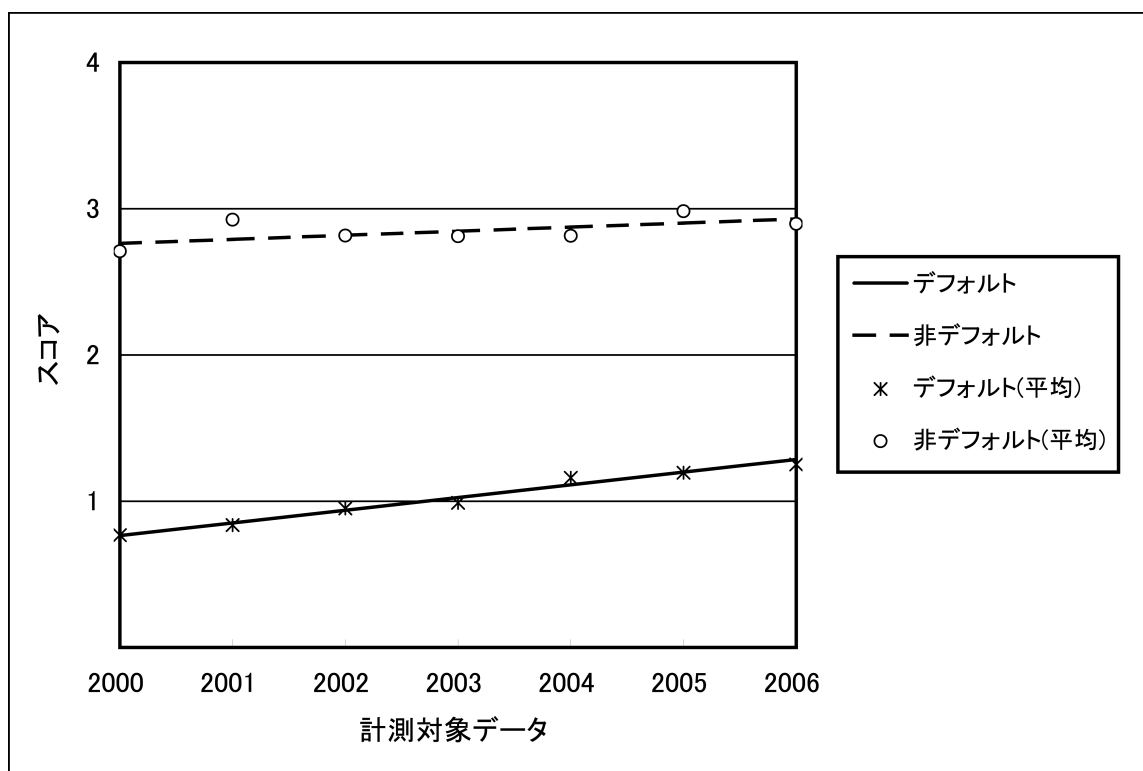


図 5: スコアの単回帰

最後に、スコアの分布の推移を定量化するためにSteel-Dwass 検定 [10][11] を行った. Steel-Dwass 検定は、複数の群がある場合に、全ての群間で Wilcoxon 検定を実施し、得られた p 値の多重性を調整する手法である. Steel-Dwass 検定は Tukey 検定のノンパラメトリック版とも言える. 検定結果を表 6 に示す.

表 6: デフォルトのスコアの年度間の Steel-Dwass 検定

	2000	2001	2002	2003	2004	2005	2006
2000	–	0.2483	0.0000	0.0000	0.0000	0.0000	0.0000
2001	3.16	–	0.0000	0.0000	0.0000	0.0000	0.0000
2002	8.79	5.44	–	0.9999	0.0000	0.0000	0.0000
2003	9.78	6.55	1.40	–	0.0000	0.0000	0.0000
2004	17.02	13.75	8.99	7.23	–	1.0000	0.1629
2005	17.39	14.30	9.81	8.12	1.26	–	0.9145
2006	16.71	14.14	10.39	9.00	3.33	2.21	–

表 6 の対角線より右上の三角形には p 値を、左下の三角形には Z 統計量の絶対値を出力している. また、検定全体での有意水準を 1%とした時に、有意差がある年度間には網掛けをしている.

年度間のデフォルトスコアの中央値に有意差のない計測対象データの組み合わせは、2000年度と2001年度、2002年度と2003年度、2004年度～2005年度までの期間、計6つとなる。他の年度の組み合わせ全てにおいては有意差がある。以上の通り、Steel-Dwass 検定の結果、全ての年度のデフォルトスコアの中央値が同じであるとの仮説は1%の有意水準で棄却される。特に、網掛け部分の組み合わせで有意差が確認されている事とこれまでの結果(図4, 図5)を考え併せるとこれらの年度でのデフォルトスコアに上昇傾向がある事が強く示唆される。

6 おわりに

本分析では、2000年度から2006年度のRDBデータベースのAR値が時系列に低下傾向にある事を示した。また、そのAR値の低下がデフォルトスコアの高スコア側へのシフトに起因する事を示した。さらに、部分的にはあるが、信用リスクモデルの性能が、モデル構築に使用したデータ年度には大きく依存しない事を確認した。

本論では分析にRDBデータベースを使用した。RDBデータベースは国内の銀行(メガバンク、地銀、信金を含む)のデータを幅広く含んでおり、日本全国の中小企業の財務データの傾向を捕らえている。その為、本分析結果から2000年度から2006年度の間日本全国のデフォルト企業の財務指標に何らかの要因が影響を与えていると考えられる。その要因として、幾つかの候補を以下に示す。

- 一点目として、各行における信用リスクモデルの使用がAR値低下の原因として考えられる。金融庁(2005)[7]に記されているように、スコアリングモデル(信用リスクモデル)を活用した融資が2002年度以降に増加している。多くの銀行で信用リスクモデルから得たスコアにより債務者判別がなされた場合、一定スコア以下を融資対象外とすると、そのような債務者がRDBデータベースに含まれなくなる。その結果、低スコア層の債務者が減少し債務者集団全体のスコアの平均値が上昇する。更に、低スコア層の債務者の比率が高いデフォルトの債務者集団にとってこの影響は大きい。故に、デフォルトの債務者集団スコアの平均値が大きく上昇し、AR値の低下が起こった。
- 二点目に、マクロ経済因子の影響が考えられる。内閣府の月例経済報告などでも示されているように、2002年2月から日本の景気は回復基調にある。その回復基調の中、日本の中小企業の財務状態が底上げされスコアが上昇した。また、高スコアの債務者より低スコアの債務者に対して、マクロ経済因子のスコアへの寄与が大きく、非デフォルトの債務者集団に比べてデフォルトの債務者集団のスコアの平均値が大きく上昇した。その結果、AR値の低下が起こった。

一点目の信用リスクモデルの使用がAR値の低下要因であった場合、未来の何時かの時点でAR値が定常状態に達し、AR値の年度間推移が安定すると考えられる。また、二点目のマクロ経済因子に関しては、景気の拡張期だけでなく後退期のスコアの推移を観察する事で、より明確な結論が得られると考察する。故に今後の課題として、継続的にRDBデータベースを調査し、スコアやAR値に影響を与えている要因の特定に努めたい。

最後に、本分析を進めるにあたり有益なコメントを頂きました一橋大学大学院 国際企業戦略研究科 三浦 良造教授に謝意を表します。

参考文献

- [1] Alfred Hamerle, Robert Rauhmeier, Daniel Rösch. *Uses and Misuses of Measures for Credit Rating Accuracy*. University of Regensburg(2003).
- [2] Bernd Engelmann, Robert Rauhmeier(Editors). *The Basel II Risk Parameters*. Springer Berlin, Heidelberg(2006).
- [3] Box, G.E.P. and Cox, D.R. *An Analysis of Transformations*. Journal of the Royal Statistics Society, B-26, 211 - 252(1964).
- [4] Donald van Deventer, Kenji Imai. *CREDIT RISK MODELS AND THE BASEL ACCORDS*. John Wiley & Sons Inc(2003).
—邦訳. 三浦 良造 (訳者代表). 信用リスクモデル入門 バーセル合意とともに. 東洋経済新報社 (2007).
- [5] Michel Grouhy, Dan Galai, Robert Mark. *Risk Management*. Mcgraw-hill(2000).
—邦訳. 三浦 良造 (訳者代表). リスクマネジメント. 共立出版 (2004).
- [6] 大久保 豊, 岸本 義之, 根本 直子, 本島 康史, 山本 真司. 銀行経営の理論と実務. 金融財政事情研究会 (2003).
- [7] 金融庁.『リレーションシップバンキングの機能強化に関するアクションプログラム』に基づく取組み実績と総括的な評価について. (2005).
- [8] 坂和 正敏, 田中 雅博. 遺伝的アルゴリズム. 朝倉書店 (1995).
- [9] 丹後 俊郎, 山岡 和枝, 高木 晴良. ロジスティック回帰分析. 朝倉書店 (1996).
- [10] 永田 靖, 吉田 道弘. 統計的多重比較法の基礎. サイエンティスト社 (1997).
- [11] 浜田 知久馬 (監修), 臨床評価研究会 基礎解析分科会 (執筆). 実用 SAS 生物統計ハンドブック. サイエンティスト社.
- [12] 安田 隆二, 大久保 豊. 信用リスク・マネジメント革命. 金融財政事情研究会 (1998).
- [13] 山下 智志, 川口 昇, 敦賀 智裕. 信用リスクモデルの評価方法に関する考察と比較. 金融庁金融研究研修センター (2003).