



リスク シミュレーター ユーザーガイド

Johnathan Mun, Ph.D., MBA, MS, CFC, CRM, FRM, MIFC

RISK SIMULATOR 2012

本書および本書内に記載されているソフトウェアは、使用許諾契約に基づいて提供され、使用承諾契約条項にしたがって使用または複製することができます。本書に記載された内容は情報の提供のみを目的としており、予告なしに変更されることがあります。リアル オプションズ バリュエーション 株式会社は、本書の内容についていかなる責任も負いません。

使用承諾契約条項で許可されている場合を除き、本書の一部または全部を リアル オプションズ バリュエーション 株式会社の事前の書面による許可なく、電子的、機械的、録音を含むいかなる手段や形式によっても、複製、検索システムへの保存、または伝送することを禁じます。

本書はリアル オプションズ バリュエーション 株式会社の創業者兼 CEO, ジョナサン・マンによる著作の使用承諾契約に基づいています。

著作権：ジョナサン マン

著作、デザインおよび印刷：アメリカ合衆国

本書の複製の購入がご希望の場合は、リアル オプションズ バリュエーション 株式会社の下記のメールかホームページを通してご連絡願います。

Admin@RealOptionsValuation.com または、 www.realoptionsvaluation.com

© 2005-2012 ジョナサン マン著作より、著作権法によって保護されています。

Microsoft® は、マイクロソフト コーポレーションの米国および他の諸国での登録商標です。その他すべての商標は、核当する各社が所有しています。

目次

1. はじめに	7
2011/2012 年版での新しいことがら	13
リスクシミュレーターの性能の包括的な一覧表.....	13
2. モンテカルロ・シミュレーション	19
モンテカルロ・シミュレーションとは?.....	19
リスク・シミュレーターを使用するにあたって.....	20
ソフトウェアの高度な概略.....	20
モンテカルロシミュレーションの実行について	21
1. 新規シミュレーションプロファイルの作成にあたって.....	21
2. 入力仮定の指定.....	24
4. シミュレーション実行.....	28
5. 予測結果の解釈.....	28
相関とエラーコントロール.....	35
相関の基礎.....	35
リスクシミュレーションで相関を適用するにあたって.....	36
モンテカルロ・シミュレーションでの相関の影響	37
予測統計を読解するにあたって.....	43
分布の中心を測定するにあたって—第一次モーメント.....	43
分布の広がり測定—第二次モーメント	43
分布の歪度を測定するにあたって—第三次モーメント.....	45
分布での破局的なテールイベントを測定するにあたって—第四次モーメント.....	46
3. 予測するにあたって	48
様々なタイプの予測法の技術.....	48
リスクシミュレーターで予測ツールを実行するにあたって.....	50
時系列分析	51
多変数回帰	55
確率的予測法	60

非線形外押法	63
Box-Jenkins ARIMA 高度な時系列.....	66
自己 ARIMA (Box-Jenkins ARIMA 高度な時系列)	72
基本的な計量経済学.....	73
J-S 曲線の予測	74
GARCH ボラティリティの予測.....	76
マルコフ連鎖 (Markov Chains).....	79
最尤推定モデル (MLE) : Logit, Probit, Tobit.....	79
スプライン (立方スプラインの補入と外押法).....	83
4. 最適化	85
最適化の方法	85
最適化と連続的な決定変数.....	88
最適化と離散的整数の変数.....	95
5. リスクシミュレーション分析ツール	101
シミュレーションでの竜巻と感度ツール.....	101
感度分析.....	108
分布的な適合: 単一変数と複数の変数.....	112
ブートストラップシミュレーション.....	116
仮説実験.....	119
データ採取とシミュレーション結果の保存.....	121
レポートの作成	122
回帰と予測の診断ツール.....	124
統計的分析ツール	134
分布的な分析ツール.....	138
リスクシミュレーター2011/2012 の新しいツール	143
乱数の生成、モンテカルロ 対 ラテン・ハイパーキューブとコンピュータ相関法...	143
非季節性データーと傾向除去データー.....	144
主成分分析	146

構造変化分析	146
傾向ラインの予測法.....	147
モデルの確認ツール.....	149
分配的なパーセンタイル適合ツール.....	150
分布グラフと表: 確率分布ツール.....	151
ROV BizStats	155
ニューラル・ネットワークと組み合わせ的ファジィ論理予測法.....	160
ゴールシークの最適化.....	164
単一変数の最適化	165
遺伝子的アルゴリズムの最適化.....	166
ROV 決定木モジュール	168
決定木.....	168
シミュレーションのモデリング.....	169
ベイジアン分析	169
完璧な情報、ミニマックス、マクシミニ分析、リスクプロファイル値と完璧ではない情報の値の予測	169
感度性	170
シナリオ表	170
ユーティリティ機能の生成.....	170
役立つ情報と技法	179
TIPS: 仮定(入力仮定ユーザーインターフェースの設定).....	179
TIPS: コピーと貼り付け.....	179
TIPS: 相関.....	180
TIPS: データー診断と統計的な分析.....	180
TIPS: 分配的な分析、グラフと確立表.....	181
TIPS: 効率的フロンティア.....	181
TIPS: 予測セル.....	181
TIPS: 予測の結果グラフ.....	181
TIPS: 予測法.....	181
TIPS: 予測法: ARIMA	182

TIPS: 予測法: 基本的な計量経済学.....	182
TIPS: 予測法: ロジット、プロビットとトビット.....	182
TIPS: 予測法: ストキャスティック過程.....	182
TIPS: 予測法: 傾向ライン	182
TIPS: 関数セル.....	183
TIPS: 練習とビデオ学習をはじめるにあたって.....	183
TIPS: ハードウェアの ID.....	183
TIPS: ラテン・ハイパーキューブ・サンプリング(LHS) 対 モンテカルロ・シミュレーション(MCS)	183
TIPS: オンラインの資源.....	184
TIPS: 最適化.....	184
TIPS: プロファイル.....	184
TIPS: 右クリックショートカットと他のショートカットキー.....	185
TIPS: 保存.....	185
TIPS: サンプリングとシミュレーション法.....	185
TIPS: ソフトウェアの開発キット(SDK)と DLL ライブラリー.....	185
TIPS: Excel でリスクシミュレーターをはじめるにあたって.....	186
TIPS: 超高速シミュレーション.....	186
TIPS: 竜巻分析.....	186
TIPS: トラブルシューター.....	187

1. はじめに

リスク・シミュレーター ソフトウェアへようこそ

リスク シミュレーター (RiskSim)の 1.1 バージョンはモンテカルロ・シミュレーションを使用した、予測および最適化の為のソフトウェアです。このソフトウェアは Microsoft .NET C#で構成されており、アドインされている Excel と一緒に活用します。このソフトウェアは互換性がある為、Real Options Valuation, Inc.のソフトウェア Real Options Super Lattice Solver (SLS)や、Employee Stock Options Valuation Toolkit (ESOV)とも使用が可能です。

このマニュアルでユーザーに役立つ十分な情報を記載することに努力しましたが、ソフトウェアの創作者（例えば、ジョナサン・マン著作のリアル オプションズ分析、第2版、ワイリー ファイナンス 2005；リスクのモデル化:モンテカルロ・シミュレーションの適用、リアル オプションズ分析、予測、そして、最適化、第2版、ワイリー 2006；そして、従業員株式購入選択権の評価 (2004 FAS 123R), ワイリー ファイナンス 2004 のトレーニング DVD、ライブトレーニングおよび本書等の代用にならないことをご承諾ください。この商品についての詳細情報はホームページでご覧になってください。
www.realoptionsvaluation.com

リスク・シミュレーターのソフトウェアには次のモジュールがあります。

- モンテカルロ・シミュレーション（パラメトリックおよびノンパラメトリックシミュレーションで42の分布確率と同時に様々なシミュレーションプロファイルの実行。また切断、相関シミュレーション、カスタマイズ シミュレーション、正確さとエラーをコントロールしたシミュレーション、その他、アルゴリズムを使用したシミュレーション等の実行）。
- 予測方法（Box-Jenkins ARIMA、重回帰分析、非線形外押法、確率過程、そして時系列分析の実行）。
- 不確定状況での最適化（シミュレーションの実行、無実行とは無関係にポートフォリオやプロジェクトの最適化の為に離散整数と連続変数を使用した最適化を実行）。
- モデル化と分析ツール（ブートストラップ・シミュレーション、仮説実験、分布適合などと同様に竜巻、スパイダー および 感度分析の実行）。

Real Options SLS のソフトウェアは単一、または複合的なオプションをコンピューター測定する他に、ユーザーのニーズに応えた調節が可能なオプションモデルを使用することも出来ます。このソフトウェアには、次のモジュールがあります。

- 単一アセット SLS（カスタマイズ オプション解決同様に、放棄、チューザー、短縮、延期や拡張のオプションなどの解決）。
- 複合資産と複合変化 SLS（継続的な複合変化の解決、根本的な資産や多段階の解決オプション、放棄、チューザー、短縮、延期や拡張等と継続的な多段階の結合、そして切り替えのオプション。ちなみにカスタマイズ オプションの解決にも使用できます）。
- 多項の SLS（三項式（平均-回帰）オプション、四項式（jump-diffusion）オプションそして、五項式（レインボー）オプションの解決）。
- Excel アドイン機能（全ての前例のオプションの解決の他に閉形式モデルが加わった解決や、Excelに基づいた環境内でのカスタマイズ オプションの解決）。

ソフトウェアのインストール

このソフトウェアをインストールするには、画面上の指導に従ってください。ちなみに最低限の要求システムが揃っていないとインストールが出来ないことをご承知ください。

- Pentium IV プロセッサーまたは、それ以上
- Windows XP または、Vista, Windows 7
- Microsoft Excel XP, 2003, 2007, 2010 または、それ以上
- Microsoft .NET Framework 2.0/3.0.
- 500MB のフリースペース
- 2GB RAM を推奨
- ソフトウェア・インストールの管理者権限

新しいコンピューターのほとんどは Microsoft .NET Framework が既にプレインストールされています。リスク・シミュレーターのインストールを実行の際に .NET Framework のインストールが必要だと言うエラーメッセージが表示された場合は、インストールを中断し、インストール CD から適切な .NET Framework をインストールしてください（言語を選び忘れないでください）。.NET Framework のインストールが終え、コンピューターをリスタートした後に、リスク・シミュレーターのソフトウェアのインストールを再度行ってください。.NET Framework 2.0 / 3.0 のどのバージョンでもコンピューター

に既にインストールされていないとリスク・シミュレーターが実行できないことをご承知願います。

10 日間のトライアル無料ライセンスのファイルがソフトウェアに付属されています。後、正規版をご購入される場合は、リアル オプションズ バリュエーション株式会社のメール admin@realoptionsvaluation.com、web サイトの www.realoptionsvaluation.com、または+1 (925) 271-4438 の電話番号にご連絡ください。本社の Web サイトでは、最新のソフトウェアのダウンロードができます。また、FAQ リンクではライセンスに関する最新アップデート情報や、インストール際に起こる問題に対する解決手段がご覧いただけます。

ライセンス

ソフトウェアの正規版をご購入になられた場合は、お客様のライセンスファイルを作成する為に、ハードウェアの ID を本社にメールしていただく必要がございます。次の項目に従ってください。

Windows XP システムで、Excel XP/2003/2007/2010 を使用の場合:

- はじめに、Excel で **リスクシミュレーター (Risk Simulator) | ライセンス (License)** をクリックし、画面に表示された 11-20 桁のコードと英数字のハードウェア ID をコピーし、admin@realoptionsvaluation.com 宛てにメールしてください。（ハードウェア ID を選択し右クリックでコピーするか、e-mail ハードウェア ID リンクをクリックしてください。）後ほど、本社からご購入になられたソフトウェアの永久ライセンス ID がお客様宛てにメールされます。この永久ライセンス ID のファイルはお客様のコンピューターのハードドライブに保存してください。そして Excel を実行した後、**リスクシミュレーター (Risk Simulator) | ライセンス (License)** をクリックし、**Install License (ライセンスのインストール)** をクリックし、永久ライセンス ID のファイルを選択すれば完了です。これで、お客様はご心配なく正規版のソフトウェアを実行することができます。このプロセスにかかる時間は一分以内です。

Windows Vista/Windows 7 システムで、Excel XP/2003/2007/2010 を使用の場合:

- はじめに、Excel 2007/2010 を Windows Vista/Windows 7 内で実行してください。
リスク・シミュレーターのメニューバーでライセンスアイコンをクリックするか、

リスクシミュレーター (**Risk Simulator**) | ライセンス (**License**) をクリックし、画面に表示された 11-20桁のコードと英数字のハードウェア ID をコピーし、admin@realoptionsvaluation.com宛てにメールしてください。(ハードウェア ID を選択し右クリックでコピーするか、e-mail ハードウェア ID リンクをクリックしてください。) 後ほど、本社からご購入になられたソフトウェアの永久ライセンス ID がお客様宛てにメールされます。この永久ライセンス ID のファイルはお客様のコンピューターのハードドライブに保存してください。その後、次のステップを行ってください。Excel を実行し、リスクシミュレーター (**Risk Simulator**) | ライセンス (**License**) または、ライセンスアイコンをクリックし、**Install License (ライセンスのインストール)** をクリックし、永久ライセンス ID のファイルを選択してください。そして、Excel を再実行したら完了です。これで、お客様はご心配なく正規版のソフトウェアを実行することができます。このプロセスにかかる時間は一分以内です。

インストールの完了後、Microsoft Excel を再実行してください。インストールプロセスが成功した場合は Excel XP/2003 のメニューバーに追加のリスク・シミュレーターが表示されます。また、Excel 2007/2010 の場合はアドイングループ内に表示され、新しいアイコンバーが Excel に表示されます。(Figure1.1 をご覧になってください。)

また、スプラッシュ スクリーンが表示され、ソフトウェアが Excel 内で機能している証拠です。(Figure1.2) Figure1.3 はリスク・シミュレーターのツールバーです。前例のアイテムが Excel に表示されていれば、ソフトウェアをすぐに実行できます。次の章ではソフトウェアの使い方が細かく記載されています。

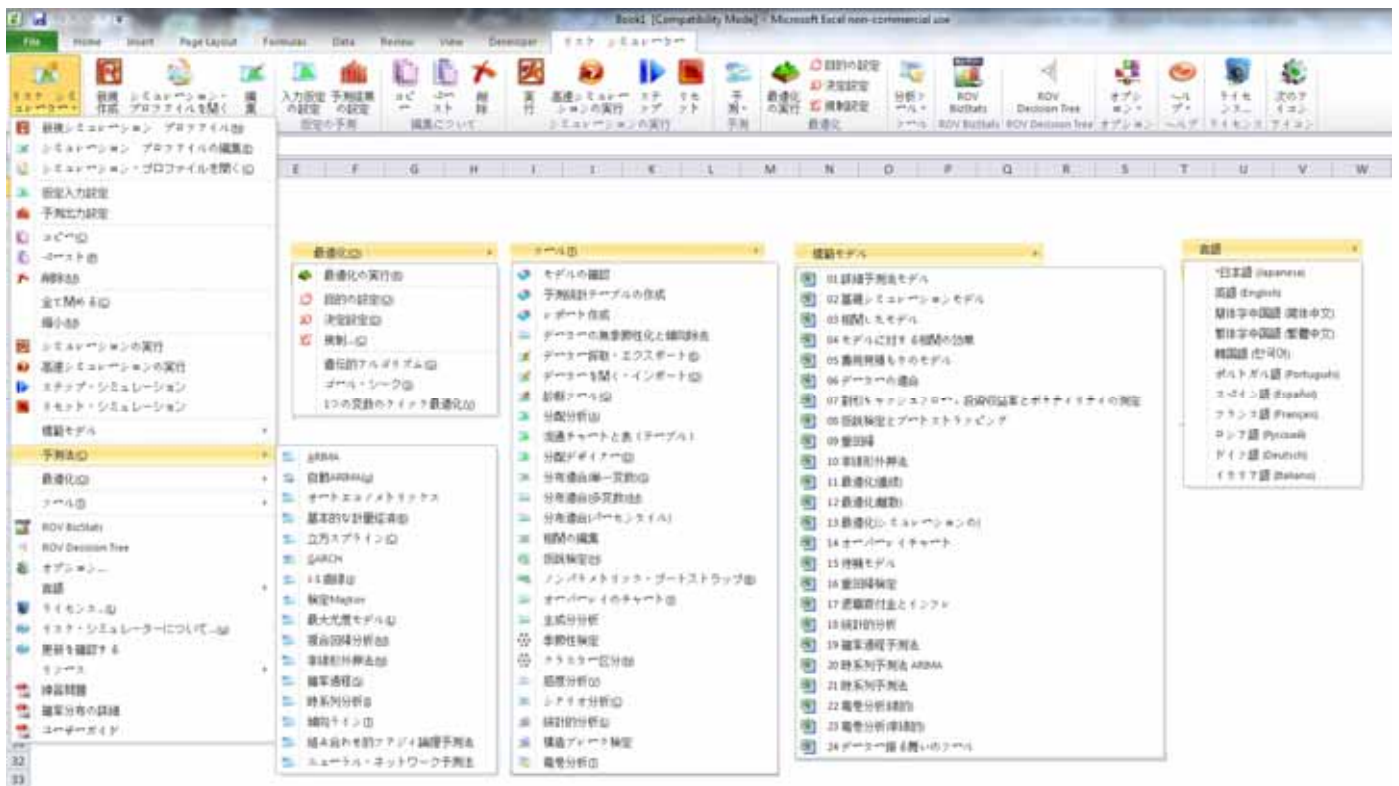


Figure 1.1 – Excel 2007/2010 のリスク・シミュレーションメニューとアイコンバー



Figure 1.2 – リスク・シミュレーター スプラッシュ スクリーン

2011/2012 年版での新しいことがら

リスクシミュレーターの性能の包括的な一覧表

下記に掲載されている目録は、リスクシミュレーターの主な性能を記述しており、強調されているアイテムは、2011/2012 年版に含まれた最新の追加を示しています。

一般的な性能

1. 11 言語で使用可能。— 英語、フランス語、ドイツ語、イタリア語、日本語、**ロシア人**、**韓国語**、ポルトガル語、スペイン語、中国語の簡体字と**繁体字**。
2. 決定木は、決定木モデルの作成と評価に使用されます。また、付加の高度な方法や分析が含まれています。
 - o 決定木モデル
 - o モンテカルロ・リスクシミュレーション
 - o 感度解析
 - o シナリオ分析
 - o ベイズ (ジョイントと事後確立の更新)
 - o 情報の予測値
 - o ミニマックス
 - o マクシミニ
 - o リスクプロファイル
3. 書庫 — 10 冊の書庫で解析理論、アプリケーションとケーススタディをバックアップ。
4. コメント・セル — セルに記述したコメントを入力、予測結果と決定変数すべてに表示、あるいは無表示に設定することが可能。
5. 細かく記述された例証モデル — リスクシミュレーターの 24 の例証モデルとモデル化のツールキットでは 300 以上のモデルが含まれています。
6. 細かく記述されたレポート — すべての解析は、細かく記述されたレポートが添付。
7. 細かく記述されたユーザーマニュアル — ステップごとに細かな説明が掲載されたユーザーマニュアル。
8. 柔軟なライセンスの取得 — 独自のリスク解析体験を味わっていただく為に、特定の機能の設定によるカスタム化が可能です。例えば、リスクシミュレーターの予測ツールにだけしか興味がない場合、予測ツールのみが使用可能なライセンスを取得する事でソフトウェアへの購入の費用を抑える事ができます。
9. 柔軟な必要条件 — Window 7 と Vista 、そして XP で起動する場合：Excel 2010、2007、2003 と互換性を持っており、MAC の場合は、バーチャル・マシーンを実行することで起動が可能。
10. 完全にカスタム化が可能な色と表 — 傾斜、3D、色、表タイプなど他！

11. 実践的な練習— リスクシミュレーターの実行のためのステップごとの解説と結果の解釈へのガイドも含まれています。
12. 複数のセルのコピーと貼り付け — 仮定、決定変数と予測のコピーと貼り付けが可能。
13. プロファイリング — シングルモデルでの複数のプロファイルの作成が可能(シミュレーションモデルの様々なシナリオを作成、複製、編集と実行ができます)。
14. Excel 2007/2010 で修正されたアイコン — より直感的でフレンドリーなアイコンのツールバーが完全に修正されました。より高い画面解像度(1280 x 760 とそれ以上)に適合する 4 つの設定が可能。
15. 右クリックのショートカット — マウスの右クリックを使用してリスクシミュレーターのツールとメニュー全てへのアクセスが可能。
16. ROV ソフトウェアの統合 —リアルオプションの SLS、モデリング・ツールキット、バーゼル・ツールキット、ROV コンパイラ、ROV 抽出と評価、ROV のモデラー、ROV の評価、ROV の最適化、ROV・ダッシュボード、ESO の評価ツールキットなど他が含まれた ROV ソフトウェアと互換性を持っています!
17. Excel での RS 関数 — Excel でマウスの右クリックで仮定と予測の設定に RS 関数を記入することが可能。
18. トラブルシューター: このツールは、システムの必要条件の確認、ハードウェアの ID の取得など、ソフトウェアを再度使用可能にしてくれます。
19. ターボスピード分析: この新しい性能は、予測と他の解説ツールの実行を猛烈的なスピード (5.2 版で強化された)で行います。解析と結果は、同じままですが、計算とレポートの生成がより速くできます。
20. ウェブリソース、ケース・スタディとビデオ — ウェブサイトから無料モデル、導入ビデオ、ケース・スタディ、ホワイトペーパーなどの素材のダウンロードが可能。

シミュレーション・モジュール

21. 6つの乱数の生成 — ROV の高度減算生成、減算乱数シャッフル生成、長期シャッフル生成、ポータブル・ランダム・シャッフル生成、クイック IEEE Hex 生成、基本最小ポータブル生成が可能。
22. 2つのサンプリング技法 — モンテカルロとラテン・ハイパーキューブ。
23. 3 コピュラ相関 — 相関されたシミュレーションへの正規コピュラ、T コピュラと擬似正規コピュラの適用。
24. 42 通りの確率分布 — 逆正弦、ベルヌイ、ベータ、ベータ 3、ベータ 4、2 項式コーシー、カイ 2 乗、余弦、カスタム、離散一様、2 重対数、アーラン、指数、指数 2、F 分布、ガンマ、幾何、グンベル Max、グンベル Min、超幾何、ラプラス、ロジスティック、対数正規 (算術) と対数正規 (対数)、対数正規 3 (算術) と 対数正規 3 (対数)、負の 2 項式、正規、パラボリック、パレート、パ

- スカル、ピアソン V、ピアソン VI、PERT、ポアソン、ベキ、ベキ 3、レイリイ、T と T2、三角、一様、ワイブル、ワイブル 3 の適用。
25. 対立的なパラメーター — パーセンタイルは、パラメーターの入力の代替法です。
 26. カスタム・ノンパラメトリック分布 — 履歴的なシミュレーションを実行し、デルファイ技法の適用によって独自の分布を作成が可能。
 27. 分布の切断 — データー境界の利用可能。
 28. Excel の関数 — Excel 内で関数を使用して仮定と予測の設定が可能。
 29. 多次元シミュレーション — 不確実な入力パラメーターのシミュレーション
 30. 確実性コントロール — 十分なシミュレーションの試行回数であるかどうかを定義します。
 31. 超高速シミュレーション — 数秒で 100,000 回、試行します。

予測法モジュール

32. ARIMA—自己回帰和分移動平均モデルの ARIMA (P,D,Q)。
33. 自己 ARIMA—最も適合するモデルを見出すために最も一般的な ARIMA の組み合わせを実行。
34. 自己計量経済学—既存するデーター（線形、非線形、相互作用、ラグ、リード、比率、相違）に最も適合するモデルの取得のために千単位のモデルの組み合わせと順列を実行。
35. 基本の計量経済学—計量経済学と線形/非線形、相互作用回帰モデル。
36. 3 次スプライン—非線形内挿法と外挿法
37. GARCH—一般化自己回帰条件付き分散不均一性モデル: GARCH, GARCH-M, TGARCH, TGARCH-M, EGARCH, EGARCH-T, GJR-GARCH, and GJR-TGARCH を使用したボラティリティの予測。
38. J 曲線—指数的な J 曲線。
39. 有限従属変数—ロジット、プロビットとトビット。
40. マルコフ・チェーン—競い合っている 2 つの要素の経時と市場有占の予測。
41. 複数の回帰—ステップワイズ法が含まれた正規の線形と非線形回帰(フォワード、バックワード、相関、フォワード-バックワード)。
42. 非線形外挿法—非線形時系列予測法。
43. S 曲線—ロジスティック S 曲線。
44. 時系列分析—レベル、傾向と季節性の予測のための 8 つの時系列分解モデル。
45. 傾向ライン—線形、非線形多項式、ベキ、対数、指数と移動平均を使用した適合度のある予測と適合。
46. ニューラル・ネットワーク予測法 (線形, ロジスティック, ハイパーボリック・タンジェント, 四編関数とハイパーボリック・タンジェント)
47. 組み合わせ的ファジ理論予測法

最適化モジュール

48. 線形最適化—一般的ニア線形の最適化と多段階の最適化。
49. 非線形最適化—ヘッセ行列やラグランジュ関数などを含めた結果の詳細。
50. 静的最適化—連続、整数と 2 値の最適化の速い実行。
51. ダイナミック最適化—最適化とシミュレーション
52. ストキャスティック最適化—2 次方程式、接線、中心、フォワード、収束基準。
53. 効率的フロンティア—多変量効率的フロンティア上のストキャスティックとダイナミック最適化の組み合わせ
54. **遺伝的アルゴリズム**—最適化問題の多様性に使用されます。
55. 重段階最適化—ローカル対グローバルの最適条件のための検定であり、最適化の実行方法のより良いコントロールと確実性の増加、そして結果の従属関係を可能とします。
56. パーセンタイルと条件的な平均値—ストキャスティック最適化のための付加の統計で、パーセンタイルの他、条件付き平均なども含まれており、リスクの測定の条件の計算にはとても重要となります。
57. **アルゴリズムの探索**—基本の単一決定変数とゴールシークのアプリケーションのためにシンプルで速く、そして有効的なアルゴリズムの探索。
58. ダイナミックとストキャスティック最適化での超高速シミュレーション—シミュレーションを超高速で実行する間に最適化を統合。

解析ツールモジュール

59. **モデルの確認**—モデルの最も多い間違えがないかを検出するテスト。
60. 関連の編集—広大な関連行列の直接な入力と編集を可能。
61. レポートの作成—モデルでの仮定と予測結果のレポート生成の自動化。
62. 統計的なレポートの作成—全ての予測統計の比較可能なレポートの生成。
63. データー診断—分散不均一性、微多数性、外れ値、自己相関、正規性、球形性、非季節性、多重共線性と相関上の検定の実行。
64. データーの抽出とエクスポート—Excel へのデーターの抽出、あるいは、テキストファイルとリスク・シムファイルのフラット化、統計レポートと予測結果のレポートの実行。
65. データーを開くとインポート—前回のシミュレーションの実行結果の回復。
66. **非季節性と傾向除去**—データーの非季節性と傾向除去。
67. 分配的な解析—42 分布すべての PDF, CDF と ICDF を計算し、確率表を生成。
68. 分配的なデザイナー—独自のカスタム分布の作成。
69. 分配的な適合(複数)—同時に複数の変数を実行。相関と相関の意味を解明。
70. 分配的な適合(単一)—連続分布上のコルモゴロフ・スミルノフとカイ 2 乗検定。レポートと分配的な仮定で完了します。
71. 仮説検定—2 つの予測が統計的に相似、あるいは相違する場合に検定します。
72. ノンパラメトリック・ブートストラップ—結果の確実性と明確性を得るための統計のシミュレーション。

73. グラフのオーバーレイ—仮定と予測結果(CDF, PDF, 2D/3D グラフタイプ)の両方の完全にカスタム化可能なオーバーレイグラフ。
74. **主成分分析**—最的な予測変数とデータアレイの減少法を検定。
75. シナリオ分析—百単位と千単位の静的 2 次元シナリオ。
76. 季節性検定—様々な季節性ラグのための検定。
77. 区分のクラスター化—データの区分のためにグループデータの統計的なクラスター化。
78. 感度性分析—ダイナミック感度性 (同時分析)。
79. **構造変化検定**—時系列データが統計的な構造変化をもっているかどうかを検定。
80. 竜巻分析—感度性、スパイダー、竜巻分析とシナリオ表の静的摂動の検定

統計的と BizStats モジュール

81. **分配的なパーセンタイル適合**—最的な適合分布を検出するためのパーセンタイルと最適化の使用。
82. **確立性分布のグラフと表**—45 通りの確率分布を全 4 次モーメント、CDF、ICDF、PDF、グラフ、分配的な複数のオーバーレイグラフ を実行し、確立性分布の表を生成します。
83. 統計的な分析—記述的な統計、分派的な適合、ヒストグラム、グラフ、非線形外挿法、正規性検定、ストキャスティックパラメーターの推定、時系列予測法、傾向ラインの予測など。
84. **ROV BIZSTATS**—130 以上のビジネス統計と分析モデル:

Absolute Values, ANOVA: Randomized Blocks Multiple Treatments, ANOVA: Single Factor Multiple Treatments, ANOVA: Two Way Analysis, ARIMA, Auto ARIMA, Autocorrelation and Partial Autocorrelation, Autoeconometrics (Detailed), Autoeconometrics (Quick), Average, Combinatorial Fuzzy Logic Forecasting, Control Chart: C, Control Chart: NP, Control Chart: P, Control Chart: R, Control Chart: U, Control Chart: X, Control Chart: XMR, Correlation, Correlation (Linear, Nonlinear), Count, Covariance, Cubic Spline, Custom Econometric Model, Data Descriptive Statistics, Deseasonalize, Difference, Distributional Fitting, Exponential J Curve, GARCH, Heteroskedasticity, Lag, Lead, Limited Dependent Variables (Logit), Limited Dependent Variables (Probit), Limited Dependent Variables (Tobit), Linear Interpolation, Linear Regression, LN, Log, Logistic S Curve, Markov Chain, Max, Median, Min, Mode, Neural Network, Nonlinear Regression, Nonparametric: Chi-Square Goodness of Fit, Nonparametric: Chi-Square Independence, Nonparametric: Chi-Square Population Variance, Nonparametric: Friedman's Test, Nonparametric: Kruskal-Wallis Test, Nonparametric: Lilliefors Test, Nonparametric: Runs Test, Nonparametric: Wilcoxon Signed-Rank (One Var), Nonparametric: Wilcoxon Signed-Rank (Two Var), Parametric: One Variable (T) Mean, Parametric: One Variable (Z) Mean, Parametric: One Variable (Z) Proportion, Parametric: Two Variable (F) Variances, Parametric: Two Variable (T) Dependent Means, Parametric: Two Variable (T) Independent Equal Variance, Parametric: Two Variable (T) Independent Unequal Variance, Parametric: Two Variable (Z) Independent Means, Parametric: Two Variable (Z) Independent Proportions, Power, Principal Component Analysis, Rank Ascending, Rank Descending, Relative LN Returns, Relative Returns, Seasonality, Segmentation Clustering, Semi-Standard Deviation (Lower), Semi-Standard Deviation (Upper), Standard 2D Area, Standard 2D Bar, Standard 2D Line, Standard 2D Point, Standard 2D Scatter, Standard 3D Area, Standard 3D Bar, Standard 3D Line, Standard 3D Point, Standard 3D Scatter, Standard Deviation (Population), Standard Deviation (Sample), Stepwise Regression (Backward), Stepwise Regression (Correlation), Stepwise Regression (Forward), Stepwise Regression (Forward-Backward), Stochastic Processes (Exponential Brownian Motion), Stochastic Processes (Geometric Brownian Motion), Stochastic Processes (Jump Diffusion), Stochastic Processes (Mean

Reversion with Jump Diffusion), Stochastic Processes (Mean Reversion), Structural Break, Sum, Time-Series Analysis (Auto), Time-Series Analysis (Double Exponential Smoothing), Time-Series Analysis (Double Moving Average), Time-Series Analysis (Holt-Winter's Additive), Time-Series Analysis (Holt-Winter's Multiplicative), Time-Series Analysis (Seasonal Additive), Time-Series Analysis (Seasonal Multiplicative), Time-Series Analysis (Single Exponential Smoothing), Time-Series Analysis (Single Moving Average), Trend Line (Difference Detrended), Trend Line (Exponential Detrended), Trend Line (Exponential), Trend Line (Linear Detrended), Trend Line (Linear), Trend Line (Logarithmic Detrended), Trend Line (Logarithmic), Trend Line (Moving Average Detrended), Trend Line (Moving Average), Trend Line (Polynomial Detrended), Trend Line (Polynomial), Trend Line (Power Detrended), Trend Line (Power), Trend Line (Rate Detrended), Trend Line (Static Mean Detrended), Trend Line (Static Median Detrended), Variance (Population), Variance (Sample), Volatility: EGARCH, Volatility: EGARCH-T, Volatility: GARCH, Volatility: GARCH-M, Volatility: GJR GARCH, Volatility: GJR TGARCH, Volatility: Log Returns Approach, Volatility: TGARCH, Volatility: TGARCH-M, Yield Curve (Bliss), and Yield Curve (Nelson-Siegel).

2. モンテカルロ・シミュレーション

モンテカルロ・シミュレーションとは、モナコの有名な賭博都市で名づけられ、とても有力な方法論と言われています。実務家にとって、シミュレーションは現実にかかる複雑で困難な問題の解決に導いてくれます。モンテカルロは数千、あるいは数百万もの確率サンプル経路に基づいて人工的な未来の予想図を描き出してくれます。会社でのアナリストにとって、大学院の高等数学を受講するだけでは、論理的でも実践的でもありません。優れたアナリストは、自由になる全ての手法を用いて、可能な限り最も容易で且つ実践的な方法で同じ結果を得る人たちのことをいいます。どのケースでも、そのモデルが正確な場合、モンテカルロ・シミュレーションは類似した答えにもっと数学的にエレガントな方法論を与えてくれます。それでは、もっと細かくモンテカルロ・シミュレーションについて、そしてどのような状況でしようするのかを見ていきましょう。

モンテカルロ・シミュレーションとは？

モンテカルロ・シミュレーションとは、予測、推定、リスクの分析等に役立つ乱数を用いた最も単純なシミュレーション法です。シミュレーションは、不確定な変数に対してユーザーが事前に決めた確率分布から値を反復的に採取したり、モデルの用の値を用いてモデルの幾つかのシナリオを計算します。全体として、各シナリオがそれぞれの予測値を用い得るようなモデルにて、これらのシナリオは関連する結果を生み出します。予測とは、ユーザーがモデルの重要な産出として定義した(たいてい、公式や関数を用いて)事象で、総収入、純利益、支出総額等を含みます。

ゴルフボールを大きなカゴから交代しながら何度も出す状況を思い浮かべてみてください。まず、カゴの形とサイズによって分布の入力仮定 **input assumption** (平均 100 及び標準偏差 10 の正規分布と **v s** 一様分布または、三角分布) が決まります。つまり、カゴによってはより底が深いものやより対称的なものがあり、その仮定によってどの程度の頻度で他のものよりも確実にボールを引き出す事ができるようになります。

反復的に引き出されるボールの数は、シミュレーションした回数 (**trials**) によって異なります。複数の関連する仮定を持つ規模の大きいモデルの場合は、大きな箱の中に小さな箱がいくつも入っていて、これらの小さな箱はそれぞれにゴルフボールのセットがあ

りあちこちに飛び出している状況を思い浮かべてください。例えば変数の間に相関関係があるとしたらそれらの箱は手を取り合っていることになり、他の箱のゴルフボールとは異なって、相互関連した箱のゴルフボールは縦一列に並んで飛び出していることになります。この相互作用でピックアップされたボール一つ、一つはシミュレーションの予測結果 (*forecast output*) として記録され、表に記されます。

リスク・シミュレーターを使用するにあたって

ソフトウェアの高度な概略

リスク・シミュレーション ソフトウェアにはモンテカルロ・シミュレーション、予測、最適化などの様々なアプリケーションが含まれています。

- ① このシミュレーション アプリケーションは Excel に基づいてモデル、予測のシミュレーション（分布結果）、適合分布の表示（最も適合した総計的分布を自動的に見出してくれます）、相関のコンピューター化（変数の間での関係維持）、感度の見分け（竜巻、感度の表の作成）等をカスタム及びノンパラメトリックシミュレーション同様に（分布の詳細とそれらのパラメーターを除いた履歴データを使用したシミュレーション）実行します。
- ② 予測アプリケーションでは、計量経済予測 Box-Jenkins ARIMA、自動時系列予測（季節性とトレンドと共に）多変量回帰（線形、非線形回帰）、非線形外挿法（曲線適合）そして、確率過程（ランダムウォーク、平均回帰、jump-diffusion 過程）を生み出すの使われます。
- ③ 最適化アプリケーションは、目的値を最大、最小する制約下で、重離散的な整数、連続、混合決定変数の最適化や、モンテカルロ・シミュレーションとともに静止・動的・確率最適化等の実行、そして線形、非線形の最適化にも使用できます。Real Options Super Lattice Solver は スタンドアロン ソフトウェアでリスク・シミュレーターを補完し、単純または複雑なリアルオプションの問題を解決してくれます。

モンテカルロシミュレーションの実行について

通常、Excel モデルでシミュレーションを実行する際、次の手順を行ってください。

1. 新規シミュレーションプロファイルを作成するか、既に保存してあるプロファイルを開いてください。
2. 入力過程の設定を適したセルで行ってください。
3. 予測結果の設定を適したセルで行ってください。
4. シミュレーションを実行してください。
5. 結果を解釈してください。

例：シミュレーションの実行例を試すために **Basic Simulation Model** というファイルを開いてください。このファイルは スタートメニューからも開くことができます。次の手順でファイルをクリックしてください。 **Start | Real Options Valuation | Risk Simulator | Examples** または **Risk Simulator | Example Models** に直接入ってください。

1. 新規シミュレーションプロファイルの作成にあたって

新規シミュレーションを実行するにはまず、新規シミュレーション プロファイルの作成が必要です。シミュレーション プロファイルはシミュレーションを実行するために必要な（仮定、予測、実行の環境設定等）手順の全ての設定が表示されます。プロファイルを作成すると複数のシミュレーションシナリオが設定できます。これは複数のユーザーが同じモデル（個人的な設定を指定し）を使用できるということです。または、一人のユーザーが同じモデルでいくつものシナリオ予測を異なった分布仮定等で検定することもできます。

- **Excel** を開き、新規、または既に保存してあるモデルを開いてください。（例証の場合は、模範モデルの基礎シミュレーションを開いてください）。
- **Risk Simulator** をクリックし、**New Simulation Profile** を選択してください。
- シミュレーションに適切なタイトル名を記入してください。（Figure 2.1 参照）

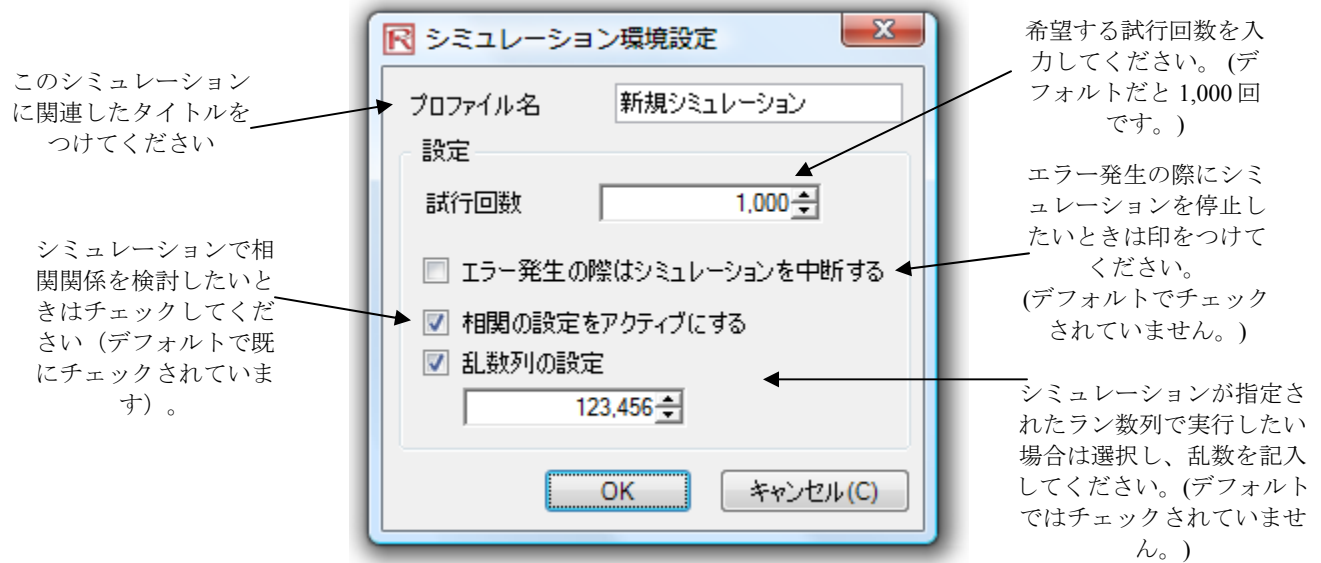


Figure 2.1 – 新規シミュレーション プロファイル

- ④ **タイトル**：シミュレーションタイトルを指定することで、一つの Excel モデルにいくつものシミュレーションプロファイルを作成することが可能になります。つまり、一つのモデルからいくつものシミュレーションシナリオを保存することが、前回使用した仮定や新規シナリオの編集を行う必要がなくなります。いつでもプロファイル名称をリスク・シミュレーター (*Risk Simulator*) | **プロファイルの編集 (Edit Profile)** を選択した後に編集することができます。
- ④ **試行数**：シミュレーションが試行される回数の指定を行うときに表示されます。つまり、1,000 試行は入力仮定に基づいた 1,000 の異なった出力結果が生成されるということです。試行回数はユーザーの希望で自由に編集することができますが、入力時に正の整数でなければいけないことを忘れないで下さい。デフォルトで既に 1,000 回と表示されます。また、精度とエラーコントロールを使用すること自動的にシミュレーションにあった必要な試行の回数が決定されます。（詳細は精度とエラーコントロールのセクションをご覧ください）。
- ④ **シミュレーションエラーの際の一時停止**：選択をすると Excel モデルでエラーが発生した際はいつもシミュレーションを停止します。例えば、モデルにエラーが生じた場合（シミュレーション試行で生成された入力値は、どれかの分散シートのセルでエラーとしてゼロで割り算されているはずでしょう）はシミュレーションが停止されます。これは Excel モデル上でコンピューターエラーがないかを確

認するために必要な項目です。したがって、モデルの構成に問題がないと予め分っている時はこのツールを使用する必要はありません。

- ② 相関のチェック：チェックすると対になっているインプット仮定間の相関が計算されます。そうでないと、相関は全てゼロと設定され、シミュレーションの実行の際、相関と入力過程の間の値は計算されません。例えば、本当に相関が存在するとすれば、相関と入力仮定の計算を行うことでもっと確実な結果が出せまし、そして負の相関があるとすれば、低い信頼度の予測が現れる傾向が見られます。この設定で相関のボックスをチェックした後、生成された各仮定の相関係数の設定ができます（詳細は相関についての章をご覧ください）。
- ③ 乱数列の指定：乱数列を指定したシミュレーションは毎回微妙に違った結果を出します。これはモンテカルロ・シミュレーションの乱数生成列の機能により、全ての乱数生成列では、理論上の事実です。しかし、プレゼンを行うとき等は既に印刷された、または表示された結果と同じ結果を表示しなければいけません。その際にはこの項目をチェックし、初期のシード数を入力してください。シード数は正の整数を使用してください。同じ初期のシード数、同じ試行数、そして同じ入力仮定を記入するとシミュレーションは常に同じ乱数列を選択し、同じ結果が出ることを保証します。

シミュレーションプロファイルが一度作成された後でも、このプロパティを編集することが出来る事を忘れないでください。これを実行する為にはまず、開いているプロファイルが編集したいプロファイルだということをご確認ください。または、[リスク・シミュレーター（Risk Simulator） | シミュレーション プロファイルの編集（Change Simulation Profile）](#) をクリックし、編集したいプロファイルを選択し、**OK** をクリックしてください。（Figure 2.2 では複数のプロファイルの中から編集したいプロファイルの選択の画面を参照しています）。ちなみにプロファイルの複製、名前の変更も可能です。

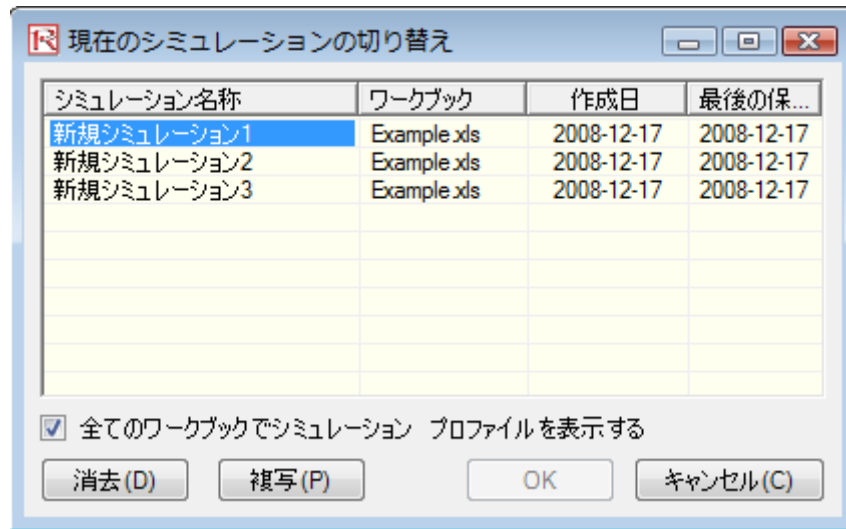


Figure 2.2 – 現在アクティブなシミュレーションのチェンジ

2. 入力仮定の指定

次の段階はモデルの入力仮定の設定です。仮定はセルに方程式や関数を除いた値を記入してください。例：モデルの入力セルには数値を、結果セルには方程式や関数を記入してください。これらの仮定と予測の呼び出しは、既にシミュレーション プロファイルが存在していないとできません。モデルに入力仮定の記入を行うには次の手順を行ってください。

- シミュレーション プロファイルがあることを確認してください。または新規プロファイルを作成するか（リスクシミュレーター **(Risk Simulator)** | **新規シミュレーション プロファイル作成 (New Simulation Profile)**）、既に保存されているプロファイルを開いてください。
- 仮定を記入したいセルを選択してください（例：基礎シミュレーションモデルのセル G8 をご覧ください）。
- **リスク シミュレーター (Risk Simulator)** | **入力仮定の設定 (Set Input Assumption)** をクリックするか、またはリスク シミュレーターのツールバーの三つ目のアイコンをクリックしてください。
- 関連のある分布法を選択し、分布のパラメーターを入力した後、モデルの入力仮定を指定する為に **OK** をクリックしてください。（Figure 2.3 参照）

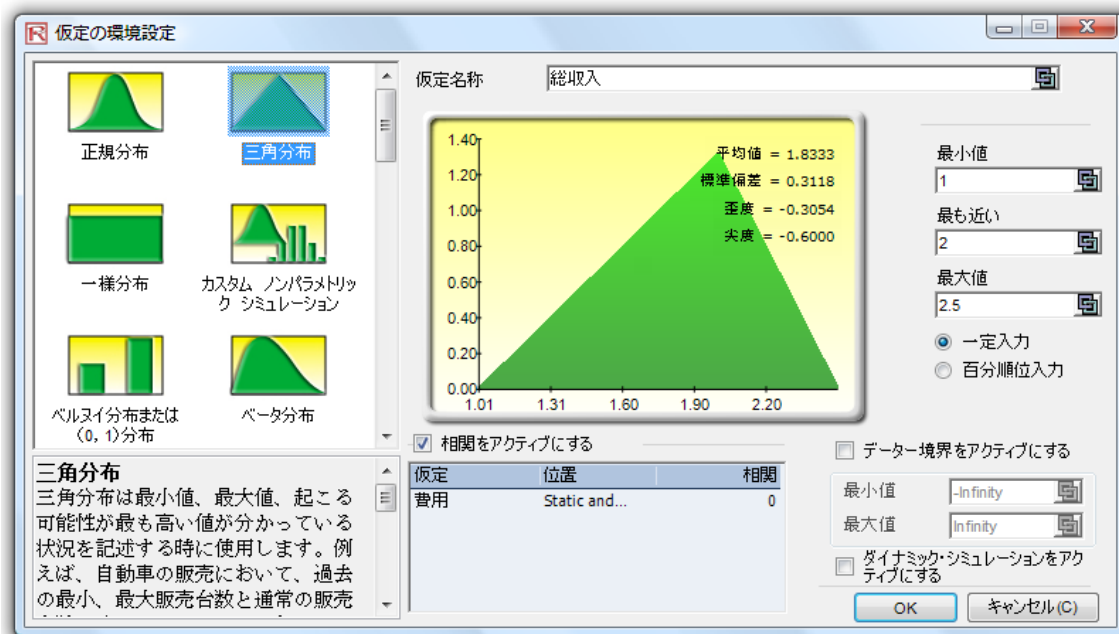


Figure 2.3 – 入力仮定の設定

仮定の環境設定にはシミュレーションにいくつかの重要な項目があります。Figure 2.4 はそれらの項目が表示されています。

- ② **仮定名称**：記入された仮定の区別がつけられるように、各仮定の名称をつけます。短くて分かりやすい名称をつけるのが一番適切でしょう。
- ② **分布ギャラリー**：左側にはソフトウェアに内蔵されている様々な分布法が表示されます。分布の画面表示を変えるには分布ギャラリーを開き大きいアイコンや小さいアイコンを右クリックしてみてください。42の分布法が表示されます。
- ② **入力パラメーター**：分布法の選択によって、関連するパラメーターが表示されます。ちなみにパラメーターは直接の入力ができますし、ワークシートの特有なセルにパラメーターをリンクさせることもできます。仮定パラメーターが変動しないと分かっている際にこのツールを使うととても便利です。パラメーターが目に見えるように表示したい時、または実行中にパラメーターを変動したいときにワークシートのセルにリンクさせると使用がより充実します。（のアイコンをクリックして入力パラメーターをワークシートセルにリンクして下さい。）
- ② **可能なデータ境界**：このツールはたいてい平均的な分析者には使用されることはありませんが、分布仮定を区分するのに使用します。例えば、正規分布を選択

したとすると、理論上、境界は負の無限と正の無限の間にあるとされます。しかし、実践ではシミュレーションされた変数は狭い範囲内で存在し、この範囲は適切な分布を区分する為に記入されます。

- ② **相関**： 対の関係にある相関は入力仮定で指定することができます。仮定が必要な場合、相関のチェックボックスをクリックすることを忘れないで下さい。このオプションは **リスク シミュレーター (Risk Simulator) | シミュレーション プロファイルの編集 (Edit Simulation Profile)** をクリックすると表示されます。相関の指定と相関がもたらすモデルへの効果の詳細は、この章の最後にある相関についての議論をご覧ください。ちなみに、分布の区分や、または他の仮定に分布を関連させることが出来ませんが、両方は同時に実行できないことをご承知ください。

- ③ **短い記述**： ギャラリーの各分布の記述が表示されます。必要な入力パラメーターと選択した分布の最適な使用法やケースが記述されています。ソフトウェアに内蔵された各分布の詳細はモンテカルロ・シミュレーションの為の確率分布の章をご覧ください。

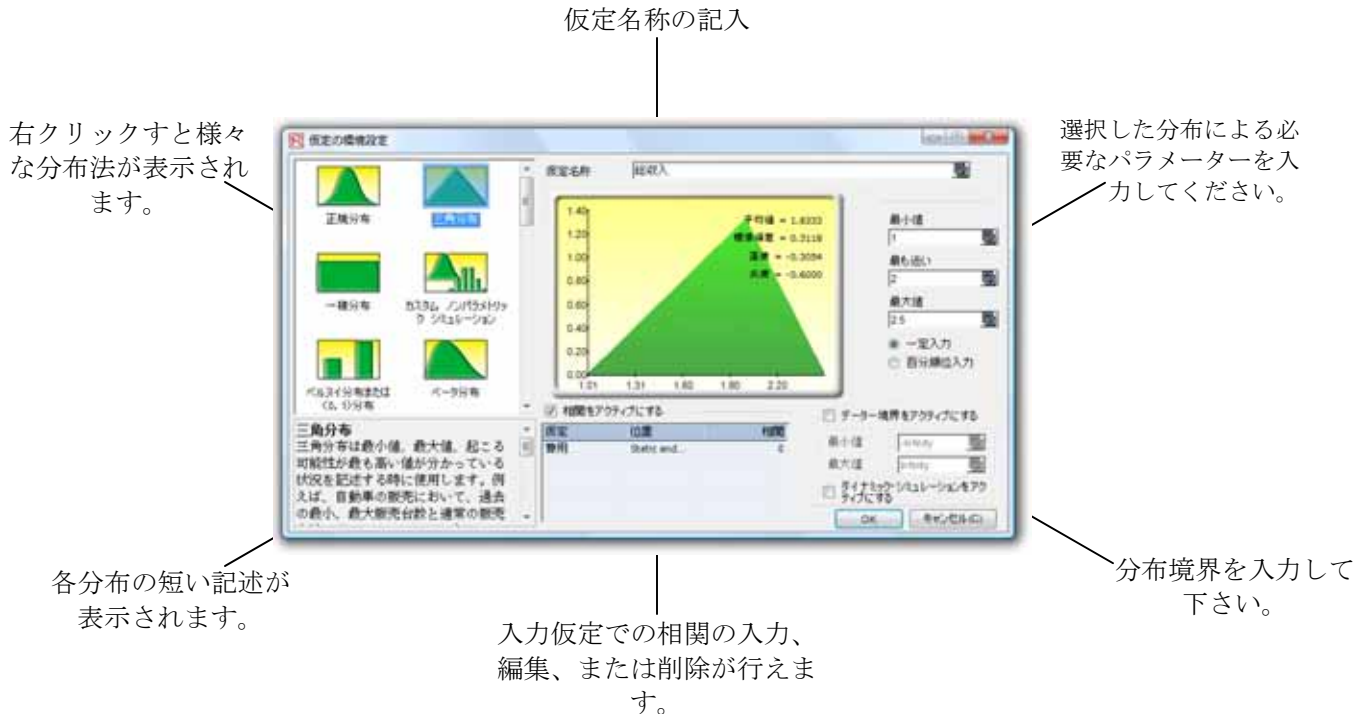


Figure 2.4 – 仮定環境設定

メモ：例証を試行している場合は他の仮定の設定を G9 で行ってください。この際に正規分布を選択し、最小値は 0.9 で最高値は 1.1 を記入して下さい。そのあと、予測結果の設定を次のステップのように行ってください。

3. 予測結果の設定

次の段階としてモデルの予測結果の設定を行います。出力セルは方程式か関数を入力してください。予測結果の設定のプロセスは次の通りです。

- 仮定を設定したいセルを選択してください（例えば、基礎シミュレーションモデル例では、セル G10）。
- **リスク・シミュレーター (Risk Simulator)** を選択し、予測結果の設定 (**Set Output Forecast**) をクリックしてください。またはリスク シミュレーターのツールバーの四つ目のアイコンをクリックしてください (Figure 1.3 参照)。
- 必要な情報を記入し、**OK** をクリックしてください。

Figure 2.5 は予測結果の環境設定を表示しています。

- ② **予測結果の名称**：予測結果のセルの名称。大きな規模のモデルで複数の予測結果のセルを設定する時に名称を区別する事で、より見やすく、より早く結果を求めることが可能になります。この単純なステップの大切さを忘れないで下さい。短くて分かりやすい名称をつけるのが一番適切でしょう。
- ② **予測の精度**：シミュレーションでどのくらいの試行数が必要かを勘で当てて見つける代わりに、精度とエラーコントロールの設定を使用して割り出す事が出来ます。シミュレーション試行回数を自動過程にすることで、エラーと精度の組み合わせのシミュレーションが計算的に必要試行回数を達成すると、シミュレーションは停止し、必要な精度の達成を連絡するので、必要なシミュレーション回数を事前に推測する必要がありません。詳細はエラーコントロールと精度の章をご覧ください。
- ② **予測結果の画面の表示**：ユーザーの希望で特定の予測画面の表示、非表示が設定できます。デフォルトで、常に予測チャートが表示されています。

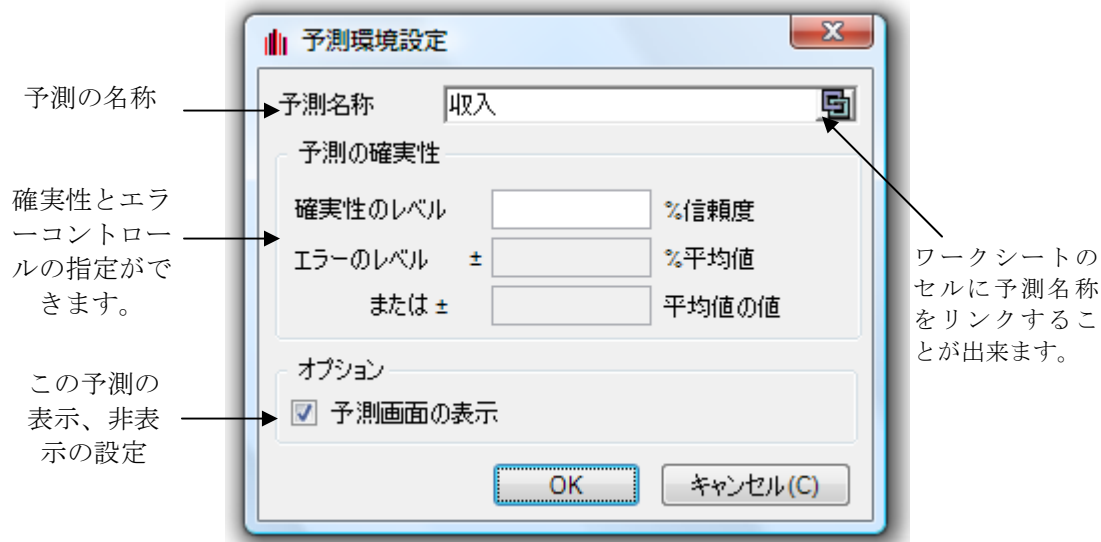


Figure 2.5 – 予測結果の設定

4. シミュレーション実行

リスク シミュレーター (Risk Simulator) | シミュレーションの実行 (Run Simulation) をクリックするか **実行 (Run)** アイコンをクリックすると (リスク シミュレーターツールバーの八つ目のアイコン) シミュレーションが実行されます。再実行する場合、または実行を中断する場合は、シミュレーションを初期化することを忘れないで下さい。この為には **リスク シミュレーター (Risk Simulator) | シミュレーションの初期化 (Reset Simulation)** をクリックするシミュレーションツールバーの 10 番目のアイコンをクリックしてください。ちなみにステップ機能 (**リスク シミュレーター (Risk Simulator) | ステップ シミュレーション (Step Simulation)** またはツールバーの九つ目のアイコン) は一度に一試行が出来ます。シミュレーションの教材として使用するには適しているでしょう。(各試行でセルに表示される仮定の値が毎回変わっていて、モデルの計算が毎行われていることが証明できます)。

5. 予測結果の解釈

モンテカルロ・シミュレーションの最終段階として予測結果のチャートの解釈があります。Figures 2.6 から 13 まで予測結果チャートとシミュレーション実行後に生成された統計が表示されます。通常、次の諸項目はシミュレーションの予測結果の解釈に必要なだとされます。

- ② 予測チャート：Figure 2.6 で参照された予測チャートはシミュレーションの総試行数から起こった値の度数の確率ヒストグラムです。縦軸は総試行数から外れて起こった特定の X 値の度数を表し、一方、累積的な度数(スムーズなライン)は予測の実行で起こった総確率の全ての値、そして X 値より下の値も表示されます。
- ③ 予測統計：Figure 2.7 の参照は予測値の分布を四つの分布段階にまとめています。これらの統計平均値の詳細は予測統計の理解の章をご覧ください。スペースバーを使うと統計図とヒストグラムの表が交代で表示されます。

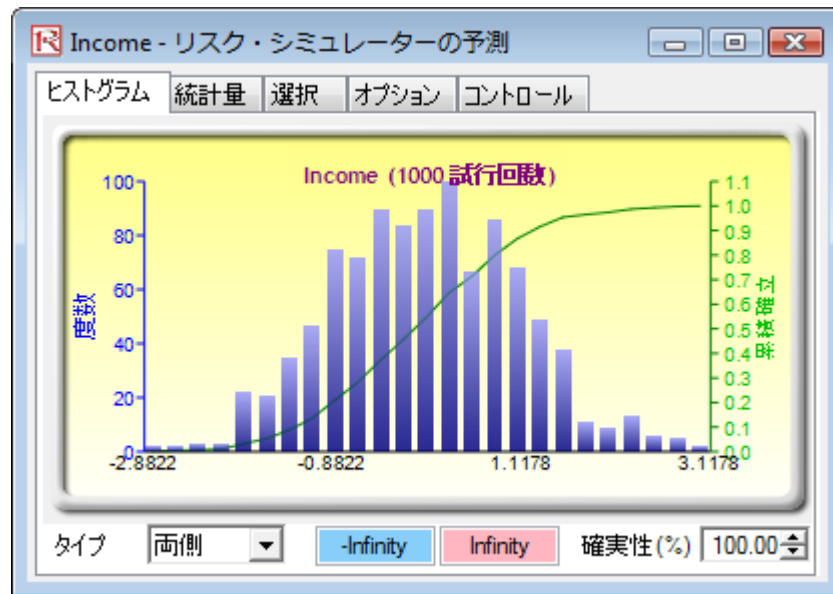


Figure 2.6 – 予測チャート

Statistics	結果
試行回数	1000
平均値	0.8626
中間値	0.8674
標準偏差	0.1933
変化	0.0374
変動系数	0.2241
最大値	1.3570
最小値	0.3019
範囲	1.0551
歪度	-0.1157
尖度	-0.4480
25分位数	0.7270
75分位数	1.0073
エラーの確実性の割合が95%の信頼度	1.3888%

Figure 2.7 – 予測統計

- ② 環境設定：予測チャートの環境設定はチャートの表示を変えてくれます。例えば、**常に前に置く (Always On Top)** を選択すると予測チャートは他のプログラムを使用していても関係なく常に表示されます。ヒストグラム解決はヒストグラムの bins 値を 5 bins から 100 bins まで変えられることが出来ます。ちなみに**データ アップデート オプション (Data Update)** はシミュレーション実行の速度 vs どのくらいの期間で予測チャートがアップデートするのかが設定できます。つまり、各試行での予測チャートのアップデートを希望すると、シミュレーションの実行速度が落ちるということです。ただし、これはユーザーの希望で設定するオプションであり、シミュレーションの結果（実行の速度以外）には影響しません。シミュレーションの実行速度の効率を上げるために、シミュレーション実行中、Excel の画面を縮小することが出来ます。画面を表示する為のメモリーが使用されなくなる為、シミュレーションの実行速度が上がるということです。**全てを削除 (Clear All)** と **全てを縮小 (Minimize All)** は開いている全ての予測チャートをコントロールします。

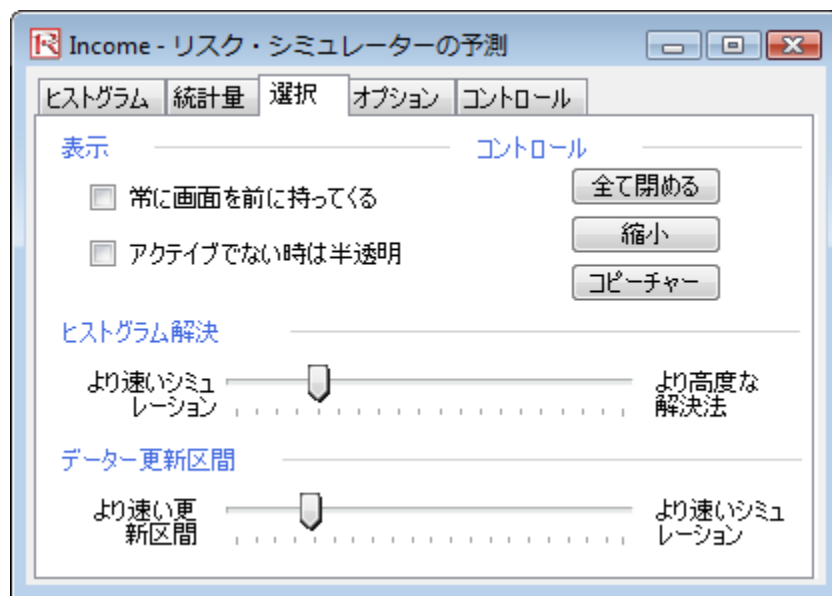


Figure 2.8 – 予測チャートの環境設定

- ③ オプション：この予測チャートオプションは全ての予測データを表示する為、または指定した区間、標準偏差に当てはまる入力・出力値をフィルターに通す為です。ちなみに精度レベルをここで指定し、統計表でエラーレベルを割り出すこ

とも出来ます。詳細は精度とエラーコントロールの章をご覧ください。次の統計を表示する（*Show the Following Statistics*）はユーザーの自由な環境設定であり平均値、中間値、最初の四分位数と4番目の四分位数系列(25番目と75番目 percentiles(百分順位))が予測チャートで表示されなければいけません。

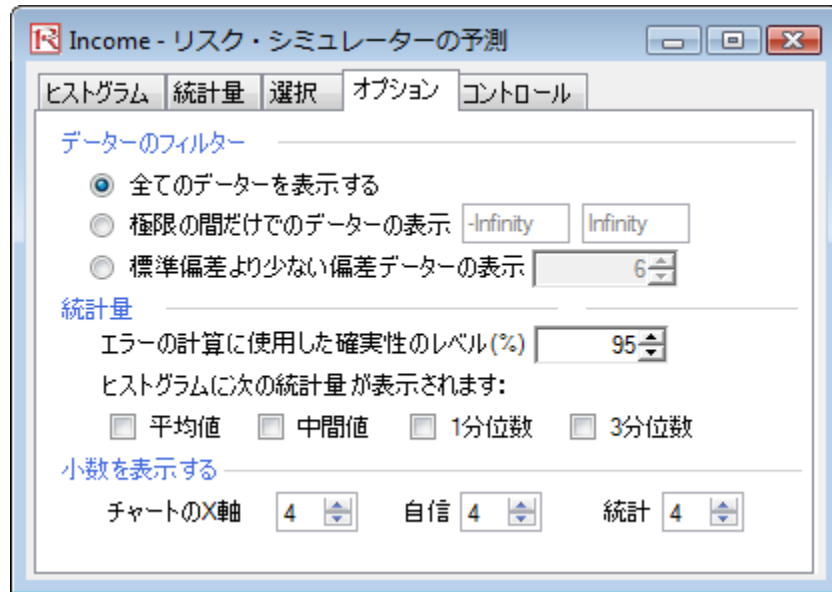


Figure 2.9 – 予測チャートオプション

予測チャートと信頼区間の使用にあたって

予測チャートでは信頼区間と呼ばれる発生確率が指定できます。つまり、二つの値が与えられたした場合、どれくらいの確率で結果値がこの二値の間に発生するのでしょうかという問題です。Figure 2.10 では 90%の確率で最終結果（このケースでは収入レベル）は\$0.5273 と \$1.1739 の区間に収まると示しています。両側信頼区間はまず、両側という形データタイプを選択し、必要な精度値を記入し（例えば：90）TAB キーを打ってください。精度値に適した二つの値が表示されます。例証には 5%の確率で収入が\$0.5273 を下回り、5%の確率で収入が\$1.1739 を上回ると表示されています。これは両側信頼区間が対称的な中間区間であるか、50 番目の percentile (百分順位)値でしょう。どちらにしろ、両方の値は同じ確率を持っています。

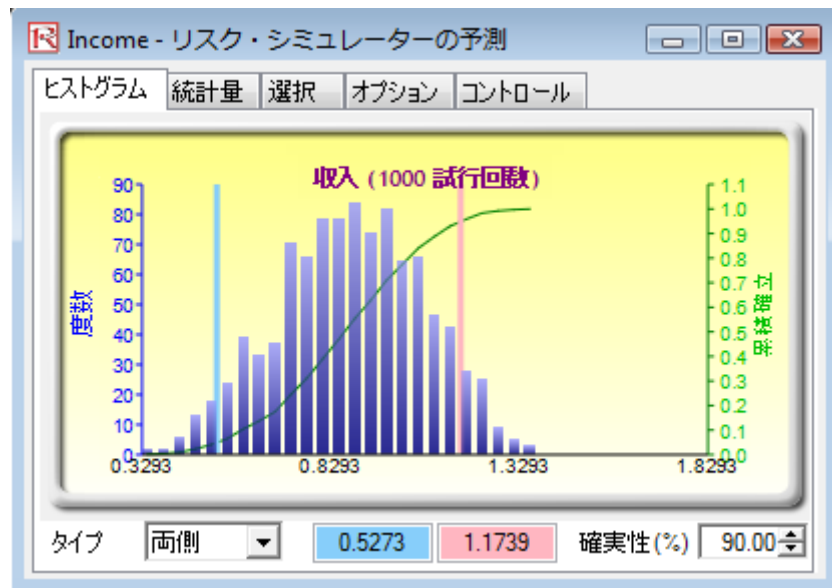


Figure 2.10 – 両側信頼区間の予測チャート

代わりに、片側確率の計算が出来ます。Figure 2.11 では左側セレクションを 95%の信頼度（例：左側を選択し、精度のレベルに 95 と記入し、TAB キーを打ってください。）とします。これは 95%の確率で収入が\$1.1739 を下回るか、5%の確率で収入が\$1.1739 を上回るということを示しています。これは Figure 2.10.で表示されている結果にしっかりと基づいて出した数値です。

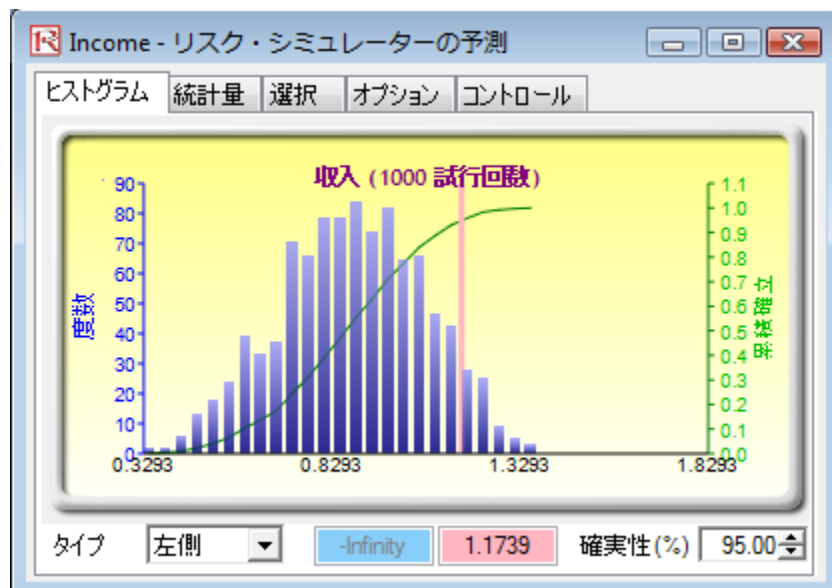


Figure 2.11 – 片側信頼区間の予測チャート

信頼区間についてももう少し付け加えると（精度レベルを与えてくれ、重要な収入値を見つけ出してくれる他）与えられた収入値の精度も計算できます。例えば、どのくらいの確率で、収入が\$1 を下回るでしょうか？という問題です。これを解決するには左側の確率タイプを選択し、1 の数値を記入し、TAB キーを打てば、適した確率値が表示されるでしょう。（この場合は 74.30%の確率で収入が\$1 を下回るという結果が出ます）。

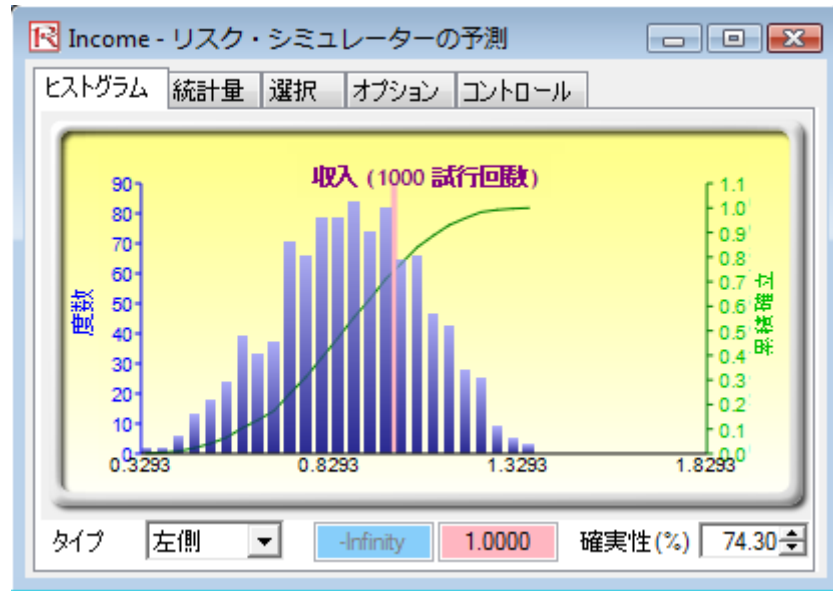


Figure 2.12 – 予測チャート確率評価

ちなみに、右側の確率タイプを選択し、1 の値を入力ボックスに記入し、TAB キーを押してください。右側の確率は値 1 を越えていることを示しています。つまり、収入が\$1 を超える確率です（このケースでは 25.70%の確率で収入が\$1 を越えていることが分かります）。

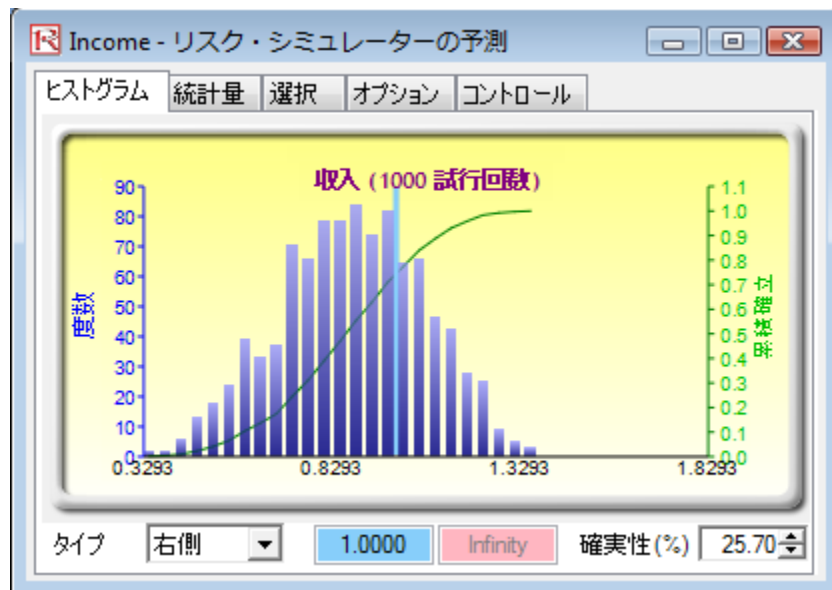


Figure 2.13 – 予測チャートの確率評価

アドバイス： 予測画面の右側の端をクリックし、ドラッグすることでサイズを変えることができます。最後にシミュレーションを再実行する前に毎回リセットすることをお勧めします。リセットには次のメニューをクリックしてください（**リスクシミュレーター (Risk Simulator) | シミュレーションのリセット (Reset Simulation)**）。確実な値や左右に値を記入した際は、チャートと結果をアップデートするのに **TAB** キーをヒットすることを忘れないでください。

相関とエラーコントロール

相関の基礎

二つの変数の間の関係の強さと方向は相関係数の測定で、-1.0 と +1.0 の間の値に当てはまります。つまり、相関係数は符号（二つの変数の間の関係が正か、負か）と、関係の大きさや強さ（高さは相関係数の絶対値、強さは関係）によって分解されます。相関係数は様々な方法で計算できます。はじめに、二つの変数 x と y を使って相関 r を手動で計算します。

$$r_{x,y} = \frac{n \sum x_i y_i - \sum x_i \sum y_i}{\sqrt{n \sum x_i^2 - (\sum x_i)^2} \sqrt{n \sum y_i^2 - (\sum y_i)^2}}$$

2 番目の方法として、Excel の CORREL 機能を使う事です。例えば、 x と y のデータポイントが A1:B10 のセルに記入されている場合、**CORREL (A1:A10, B1:B10)**と表示します。

3 番目の方法として、リスク・シミュレーターのマルチ・フィットツール (**Multi-Fit Tool**) を実行すると、結果として起こった相関マトリックスが計算され、表示されます。

相関関係は因果関係を意味しない事に注目してください。二つの無関係な確率変数は、ある相関関係を指定しなければいけませんが、原因は含意されません（例えば、太陽の黒点の活動株式市場におけるイベントは相関づけられていますが、2 つの間には、因果関係が全くありません）。

二つの一般的タイプの相関があります：パラメトリックとノンパラメトリック相関です。ピアソン相関係数法は、最も一般的な相関測定法であり、通常、単に相関係数として表示しています。しかし、ピアソン相関法はパラメトリック法であることから、相関関係にある互いの変数は基本的に正規分布で、変数の間の関係は線形でなければいけないことを示しています。モンテカルロ・シミュレーションでよく見られる、これらの条件が満たされないケースでは、ノンパラメトリックの対応側がもっと大切になります。スピアマンの順位相関とケンドールのタウの二つが代替案です。

スピアマン相関は最も一般的に使用され、モンテカルロ・シミュレーションの環境に適用するのが最も適切です。線形性または正規分布上に依存しない為、異なった分布を持った異なった変数の間の相関が適用できることを示しています。スピアマンの相関を計算する為には、まずは、全ての x と y の変数を順位別にランクしてからピアソンの相関の計算を行います。

リスクシミュレーターの場合、使用された相関は、より頑強なノンパラメトリックスピアマンのランクの相関です。但し、シミュレーションの過程を簡素化し、Excel の相関機能と一貫性を保つ為には、必要とされた相関の入力は、ピアソンの相関係数です。その後、リスクシミュレーターは、独特のアルゴリズムを適用し、過程の簡素化によって、これらをスピアマンのランク相関に変換します。ただし、ユーザーインターフェースを簡素化するには、より一般的なピアソンの積率相関の追加入力を可能にします（例、EXCEL の CORREL 機能を使用した計算）。他方、数学的なコードでは、これらのシンプルな相関を分布シミュレーションの為に基づいたスピアマンの順位相関に変換します。

リスクシミュレーションで相関を適用するにあたって

相関はリスクシミュレーターに様々な方法で適用できます:

- ① 仮定を定義する時(リスクシミュレーター(**Risk Simulator**) | **入力仮定の設定(Set Input Assumption)**))、分布ギャラリーで相関マトリックス格子内に相関を記入してください。
- ② データの存在がある時、マルチ・適合ツールで(リスクシミュレーター(**Risk Simulator**) | **ツール(Tools)** | **分布適合(Distributional Fitting)** | **複数の変数(Multiple Variables)**))、分布適合を表示し、ペアワイズ変数の間の相関マトリックスを得る為に実行します。もし、シミュレーション プロファイルが存在するのであれば、適合した仮定は自動的に関連する有相関の値を含んでいると考えられます。
- ③ データの存在がある時、リスクシミュレーター(**Risk Simulator**) | **相関の編集(Edit Correlations)**をクリックし、一人のユーザーのインターフェース内の全ての仮定のペアワイズ相関の記入を行います。

相関マトリックスが正定値行列でなければいけないことに注目してください。これは、相関は数学的に有効でなければいけないことを示しています。例えば、三つの変数の相関を計算しようとしているとします。ある特定の年の大学院生学年、彼らが週に消費するビールの数と、彼らが週に勉強する時間数。次のような相関関係が存在を仮定できます。

学年とビール: - ビールの消費が多いほど、学年が低い (試験に欠席)

学年と勉強時間: + 勉強時間が多いほど、高学年です。

ビールと勉強時間: - ビールの消費が多いほど、勉強時間が少ない (常に遊び、飲み放題)

但し、学年と勉強の間に負の相関関係を入力し、相関係数が高い場合、相関マトリックスは非正定値行列になり得ます。これは、論理に、相関の必要条件に、そして数学的な行列に対立する事になります。但し、小さい相関は時により、論理が不十分な場合でも依然として機能します。非正定値行列の、あるいは誤った相関行列を記入した際には、リスクシミュレーターは自動的に表示し、相関関係 (同一の相対的強さと同一の符号) の全体的な構成を維持しながらも、これらの相関を半正定値行列に近づける為に調節します。

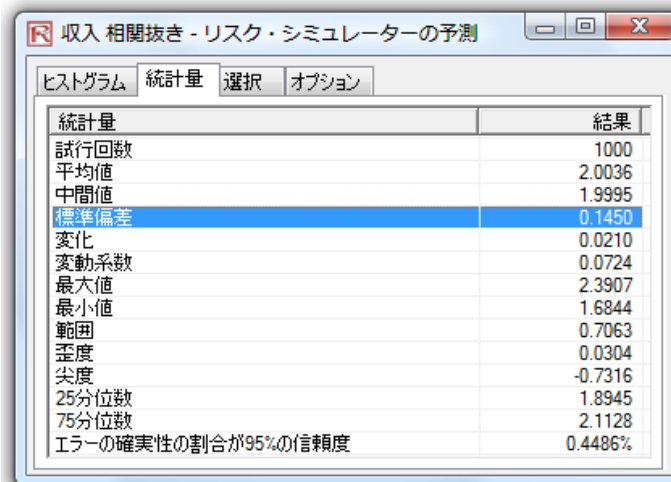
モンテカルロ・シミュレーションでの相関の影響

シミュレーションでの変数の相関に必要とされてる計算は複雑でも、結果として表示される影響はかなり明確です。Figure 2.14 は、単純相関モデルを (例証フォルダー内の相関影響モデル)を表示します。収入の為の計算は、単に価格を量でかけたものです。価格と量の間の無の相関、正の相関(+0.9)と負の相関(-0.9)で同じモデルが複製されています。

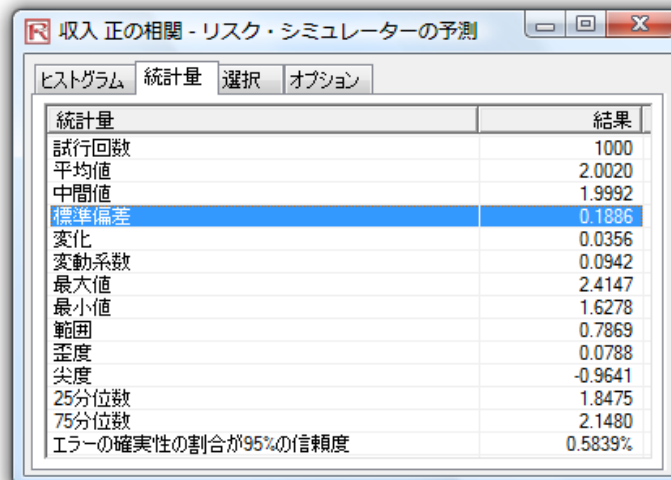
	Correlation Model		
	Without Correlation	Positive Correlation	Negative Correlation
Price	\$2.00	\$2.00	\$2.00
Quantity	1.00	1.00	1.00
Revenue	\$2.00	\$2.00	\$2.00

Figure 2.14 – 単一相関モデル

結果としての統計は、Figure 2.15 で表示されています。相関が含まれないモデルの標準偏差は、0.1886 の正の相関と 0.0171 の負の相関に比べて 0.1450 であることに注目してください。これは、シンプルなモデルでは、負の相関は分布の平均的散らばりを減少する傾向を見せ、より大きな平均的散らばりを持つ正の相関に比べ、緊密に集中した予測分布を作成する傾向があります。但し、平均値は比較的安定して残ります。これは、相関はプロジェクトの期待値をはほとんど変化させないが、プロジェクトのリスクを減少したり、増加させたりすることを意味します。



統計量	結果
試行回数	1000
平均値	2.0036
中間値	1.9995
標準偏差	0.1450
変化	0.0210
変動系数	0.0724
最大値	2.3907
最小値	1.6844
範囲	0.7063
歪度	0.0304
尖度	-0.7316
25分位数	1.8945
75分位数	2.1128
エラーの確実性の割合が95%の信頼度	0.4486%



統計量	結果
試行回数	1000
平均値	2.0020
中間値	1.9992
標準偏差	0.1886
変化	0.0356
変動系数	0.0942
最大値	2.4147
最小値	1.6278
範囲	0.7869
歪度	0.0788
尖度	-0.9641
25分位数	1.8475
75分位数	2.1480
エラーの確実性の割合が95%の信頼度	0.5839%

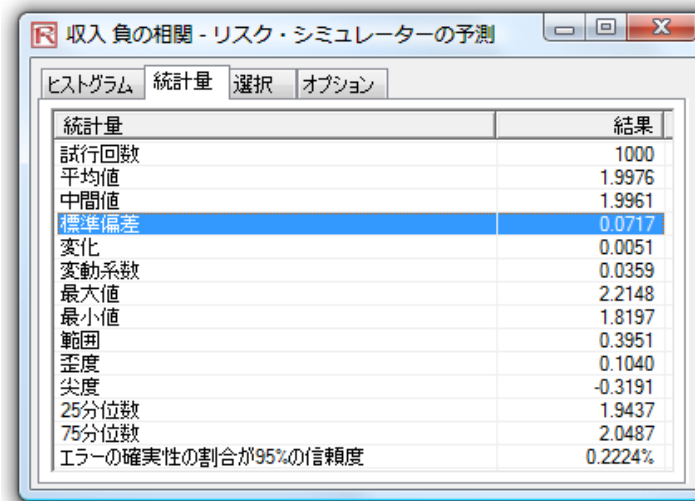


Figure 2.15 –相関の結果

Figure 2.16 は、シミュレーションの実行後に表示される結果であり、仮定の未加工データを抽出し、変数の間の相関を計算します。表は、入力仮定がシミュレーションで再現されていることを表示しています。すなわち、自身が+0.9 と-0.9 の相関を記入すれば、結果として表示されたシミュレーションの値は相関と同じになっています。

下記はシミュレーションから採取した未加工値で、これらの値は、本当に当モデルに記入された仮定の相関かどうかを確認するために相関します。ピアソン相関係数は線形パラメトリック相関で、得た結果は、記入された相関（+0.80 と -0.80）は本当に変数の間の相関を表示します。詳細はジョナサン・マン著の"Modeling Risk"(Wiley 2006)をご覧ください。

量		量	
価格	正の相関のランク	価格	負の相関のランク
1.95	0.91	1.89	1.06
1.92	0.95	1.98	1.05
2.02	1.04	1.89	1.09
2.04	1.03	1.88	1.04
1.89	0.91	1.96	0.93
1.98	1.05	2.02	0.93
2.05	1.03	2.00	1.02
1.87	0.91	1.86	1.04
1.84	0.91	1.96	1.02
2.06	1.03	1.90	1.02
1.98	1.01	1.92	1.10

Figure 2.16 –相関再現

精度とエラーコントロール

モンテカルロ・シミュレーションの1つの強力なツールは精度コントロールです。例えば、複雑なモデルで実行するのにどのくらいの試行数が十分でしょうか。精度コントロールは予め設定した精度水準に到達した場合シミュレーションを停止を可能にすることで、適切な試行数の推定という作業を取り除いています。

精度のコントロール機能は、予測をどれだけ確実にしたいかについて自らの設定を可能にしてくれます。一般的に、試行数が多いほど、信頼区間が狭くなり、統計はより確実となります。リスクシミュレーターで適用されている精度コントロールは、信頼区間の特性を使用し、いつ統計のある特定の精度が達成されるかを見出してくれます。各予測の確率精度のレベルに特定の信頼区間が指定できます。

三つの異なった言葉を混同しないように注意してください: エラー、精度と信頼です。これらが類似しているようですが、それぞれが持っているコンセプトは明らかに異なっています。これらの見本は次に表示されています。例えば、タコスシェル（タコスの皮）シェル製造業者だとし、100 のシェルが入った箱の中の平均どれだけのシェルが割れているかが知りたいとします。一つの方法としては、100 のシェルが入ったプリパッケージされた箱のサンプルを取り寄せ、開き、どれだけのシェルが不良かを数えます。一日に 1,000,000 箱（これが母集団となります）を製造し、ランダムでその中から 10 箱を開けます（これがサンプルサイズとなり、シミュレーションでの試行数とも考えられます）。開けた箱の中から不良のシェルの数は次の通りです: 24, 22, 4, 15, 33, 32, 4, 1, 45, と 2。計算された不良品のシェルの平均数は 18.2 です。これらのサンプル、または 10 の試行回数に基づいた平均値は 18.2 ユニットなのに対して、サンプルに基づく 80% の信頼区間は 2 から 33 ユニットになります（これは、実行された試行回数、またはサンプルサイズに基づいた 80%において、結果として不良品の数は 2 から 33 ユニットの中に当てはまると言う事です）。但し、どのようにして 18.2 が正しい平均値だと定めるのでしょうか? 10 の試行回数は、結果の精度を定めるのに十分でしょうか? 2 から 33 の間の信頼区間では、広すぎ、変動的といえます。例えば 90%の頻度データコスシェルのエラーの度数が ± 2 というような一層正確な平均値が必要な場合を仮定すれば、これは、一日に製造された全ての 1,000,000 箱を開けると、これらの内 900,000 箱には、特定の平均値から平均 ± 2 の個数の範囲の不良のシェルが入っているという事を意味します。これほどの精度を持った結果を得るにはどれくらいの試行回数、またはサンプルサイズが必要となるのでしょうか? ここでは、二つのタコスがエラーレベルを示し、90%が精度を示しています。十分な試行回数が実行された場合、90%の信頼区間は、90%の精度レベルと同じになり、90%の頻度、エラー（誤差）、また信頼度が ± 2 のように一層

正確な平均の尺度が得られます。例証のように、平均値が 20 ユニットだとし、90%の信頼区間は 18 から 22 ユニットの間にあります。そしてこの信頼度 90%の精度を持ち、1,000,000 箱を開けた場合、この内の 900,000 箱には 18 から 22 ユニットの不良なタコスシェルが入っている事になります。この精度に十分な試行回数を得るには、次のサンプリングエラーの方式 $\bar{x} \pm Z \frac{s}{\sqrt{n}}$ に基づいていなければいけません。 $Z \frac{s}{\sqrt{n}}$ はタコス 2 個のエラーの度数を示し、 $\bar{x}$ は、サンプルの平均値、 Z は、90%の精度レベルの標準正規 Z-スコアものとなります。また s は、標準偏差を表し、 n は、指定された精度に対比したエラーのレベルを得る為に必要な試行回数となります。Figures 2.17 と 2.18 は、リスクシミュレーターで、精度コントロールがどの様にして、複数の予測に適用できるかを表示しています。これは、ユーザーがシミュレーションを実行するのにどれほどの試行回数が必要かを定める過程をシンプルにしてくれます。

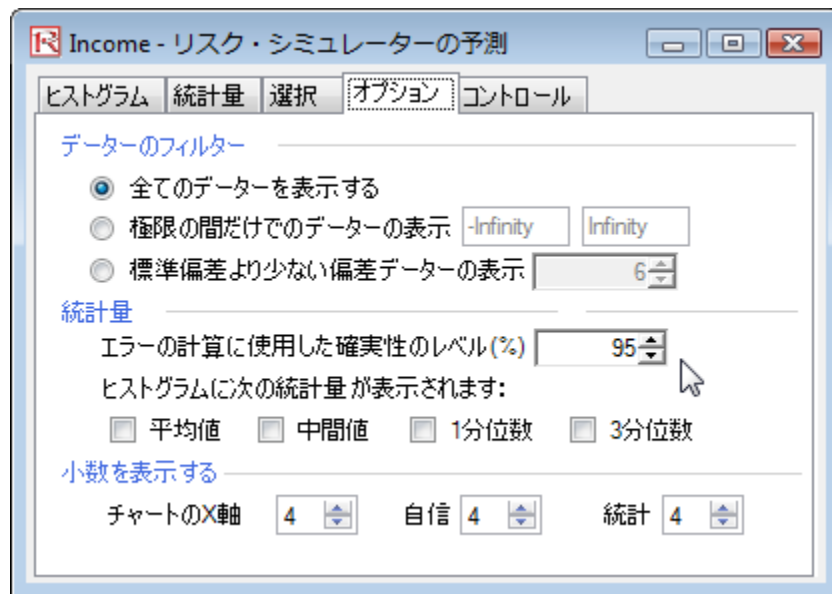


Figure 2.17 – 予測の精度レベルの設定

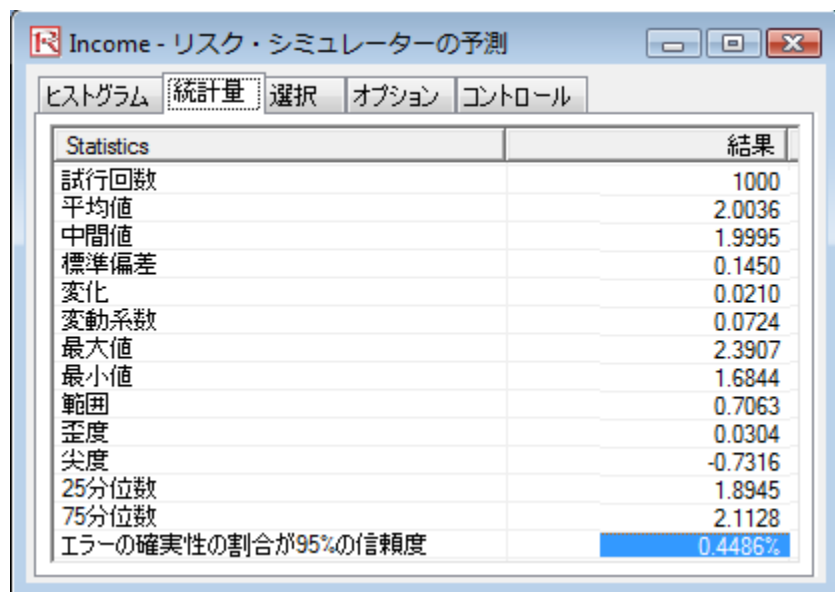


Figure 2.18 – エラーの計算

予測統計を読解するにあたって

ほとんどの分布は 4 つのモーメントまで定義可能です。第一次モーメントは、分布の位置、または中心の傾向（期待リターン）を表示し、第二次モーメントは、これらの幅、または広がり（リスク）を示し、第三次モーメントは歪みの方向（最も確率の高いイベント）を示し、第 4 次モーメントは、尖度、または テールの厚さ（破局的な損失、または利益）を示しています。これらの全ての 4 次モーメントは、分析の下でプロジェクトの最も分かりやすい視点を得る為に、実行の際に計算され、解釈されなければいけません。リスクシミュレーターは、これらの全ての 4 次モーメントの結果を予測チャートの統計の表示で与えてくれます。

分布の中心を測定するにあたって—第一次モーメント

分布の第一次モーメントは、ある特定のプロジェクトのリターンの予想の比率を測定します。つまり、プロジェクトのシナリオの位置を測定し、平均値としての確率として考えられる結果を測定します。第一次モーメントの為に最も一般的な統計は、平均（平均値）、メディアン（分布の中心）とモード（最も一般的に発生する値）が含まれています。Figure 2.19 では、第一次モーメントが表示されており、このケースでの分布の第一次モーメントは、平均、または平均値 (μ) で測定されています。

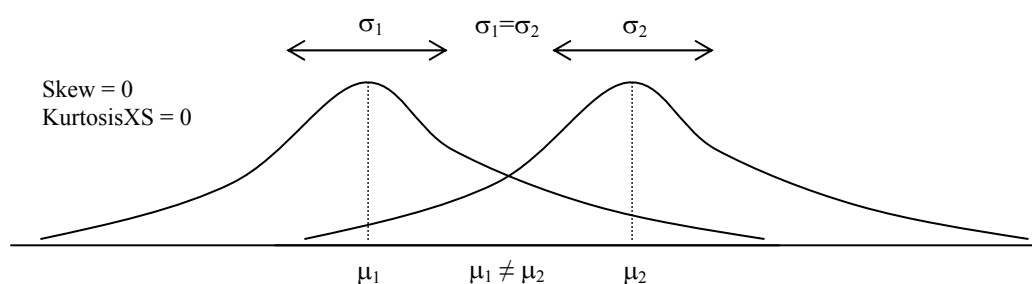


Figure 2.19 – 第一次モーメント

分布の広がり測定—第二次モーメント

第二次モーメントは、分布の広がり、つまりリスクを測定します。分布の広がり、または分布の幅は、変数の変動を測定します。これは、分布の様々な範囲に変数が落ちる

潜在性で、つまり結果のシナリオの潜在性です。Figure 2.20 は、同じ第一次モーメント（同じ平均）を持った二つの分布を表示していますが、まったく異なった第二次モーメント、またはリスクを表示しています。視覚的には Figure 2.21 でもっと明確になります。例証のように、二つの株式があるとし、小さい変動の最初の株式の変動（黒線で表示されています）に対して、二つ目の大きな価格の株式の変動（点線で表示されています）とで比較されています。一層リスクの高い株式の結果は比較的のリスクの低い株式に比べて知られていないため、明らかに、投資家は変動の大きな株式を一層高リスクと見る傾向があります。Figure 2.21 の縦軸は、株式の価格を測定しますが、リスクの高い株式ほど、潜在的な結果の幅広い範囲を示します。この範囲は、分布の幅（横軸）として Figure 2.20 で解釈されており、幅広い分布はリスクの高い資産を示しています。したがって、分布の幅、または広がりの変数のリスクを測定します。

Figure 2.20 の両方の分布は、同一の第一次モーメント、または中心の傾向を持っていますが、明らかに異なった分布であることに注目してください。この違いは、分布の幅で測定できます。数学的、そして統計的には、変数の幅、またはリスクは、範囲、標準偏差(σ)、分散、変動の係数と百分位数を含めた様々な統計を通して測定が出来ます。

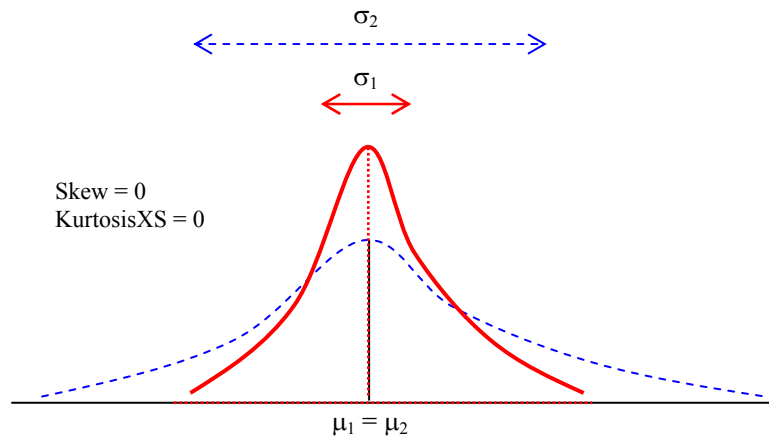


Figure 2.20 – 第二次モーメント

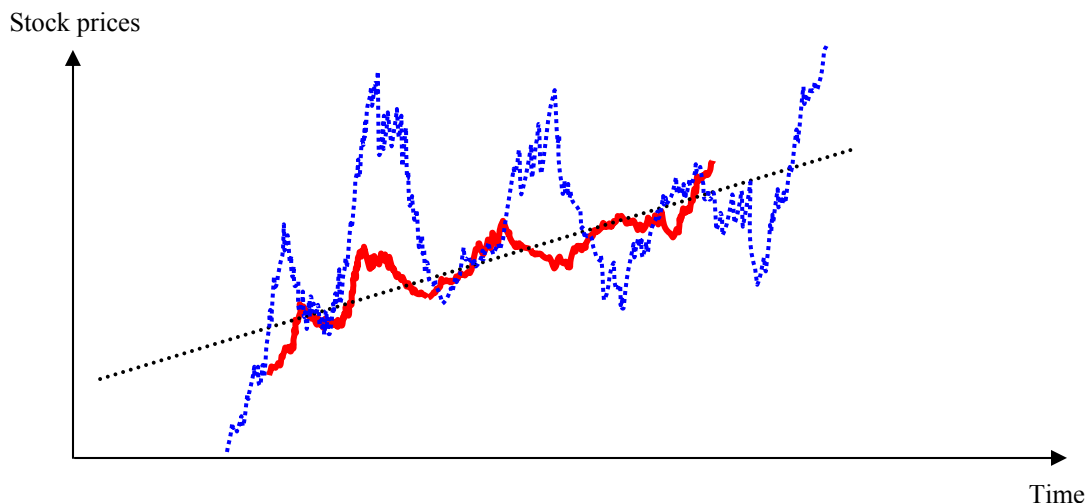


Figure 2.21 – 株価の変動

分布の歪度を測定するにあたって—第三次モーメント

第三次モーメントは、どのようにして分布が片側、または逆方面に押されるか、分布の歪度を測定します。Figure 2.22 は、負、および左側の歪み（分布のポイントのテールを左側に）を示しています。また、Figure 2.23 は、正、および右側の歪み（分布のポイントのテールを右側に）を示しています。平均は、常に分布のテールの方に歪んでいます。この別の見方として、平均は動くが、標準偏差、分散、または、幅はもの一定に留まっているということになります。もしも第三次モーメントを考慮しないで、期待リターン（例、中間、および平均）とリスク（標準偏差）のみを観察した場合、正の歪みのあるプロジェクトが誤って選択されてしまいます。例えば、横軸が、プロジェクトの純収入を表示しているとした場合、明らかに負、および左側に歪んだ分布は高確率で高いリターン(Figure 2.22)であり、高確率で低いリターン(Figure 2.23)に比較して好まれるはずです。したがって、歪んだ分布には、メディアン（中位）はリターンのより良い測定尺度であり、Figures 2.22 と 2.23 の両方ではメディアンが同じで、リスクも同一なので、純利益が負の歪みを持った分布のプロジェクトを選択するのが望ましくなります。プロジェクトの分布の歪みの考慮の失敗は、正しくないプロジェクトを選択する事を示しています（例えば、二つのプロジェクトが同じ第一次、第二次モーメントを持っている場合、これは両方とも同じリターン

とリスクのプロファイルを持っている事になるが、これらの分布の歪みはまったく異なっているかもしれません。

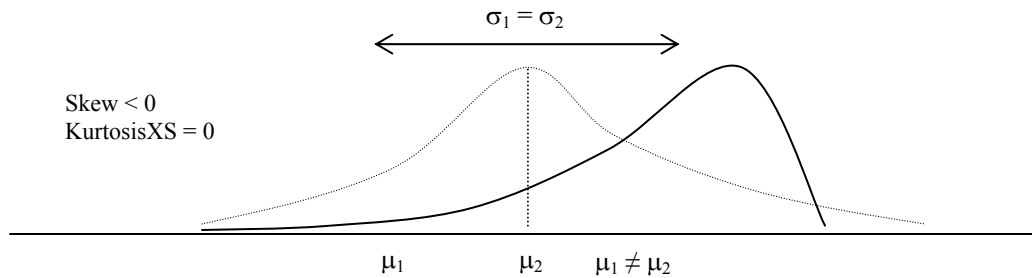


Figure 2.22 – 第三次モーメント (左側の歪み)

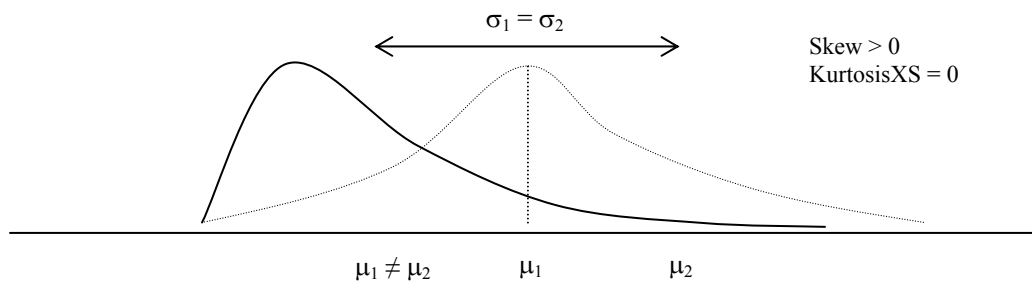


Figure 2.23 – 第三次モーメント (右側の歪み)

分布での破局的なテールイベントを測定するにあたって—第四次モーメント

第四次モーメント、または尖度は、分布のピークを測定します。Figure 2.24 は、この効果を表示しています。バックグラウンド(点線で表示されています)は、3.0 の尖度あるいは 0.0 の超過分の尖度(KurtosisXS)の正規分布です。リスクシミュレーターの結果は、尖度の平常レベルとして 0 をを用いた KurtosisXS の値を表示します。これは、正の値が尖ったテール（スチューデントの T、あるいは対数正規分布の様な急尖的分布）を示しているが、負の KurtosisXS は、平たいテール（一様分布の様な緩尖的分布）を示している事になります。太線で描かれた分布は高い尖度を持っている為、曲線の下範囲はテールに対してより厚く、中心部の範囲はより薄くなっています。この条件は、Figure 2.24 にある二つの分布のようにリスク分析により大きな影響を与え、初めの第三次モーメント（平均、標準偏差と歪度）が同一であったとしても、第四次モーメント（尖度）は異なります。この条件は、リターンとリスクが同じである場合でも、極端、および破局的なイベント（大きな

利益、または大きな損失)の発生確率は、高い尖度の分布(例、市場の株式のリターンは急尖的、または高い尖度)の方が大きい事を示しています。プロジェクトの尖度を無視することは有害となってしまいます。大抵、超過した尖度の値は、ダウンサイドのリスクは高い事を示しています(例、プロジェクトのヴァリューアットリスクの値は有意でなければいけません)。

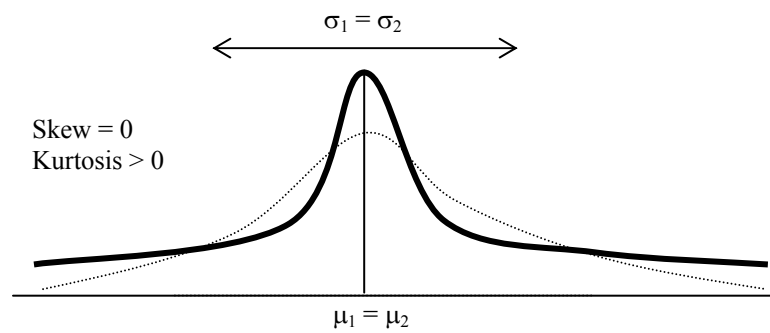


Figure 2.24 – 第四次モーメント

3. 予測するにあたって

予測とは、未来の予言をする行動にあたり、履歴的データ、また、履歴的データが存在しない場合は、未来の推測に基づいて行われます。履歴が実在する場合、定量的あるいは統計的アプローチが最適ですが、履歴的データがない場合、可能性としては定性的または判断の適用が適切となり、たった一つの方法となります。Figure 3.1 は、最も一般的な予測方法を表示しています。

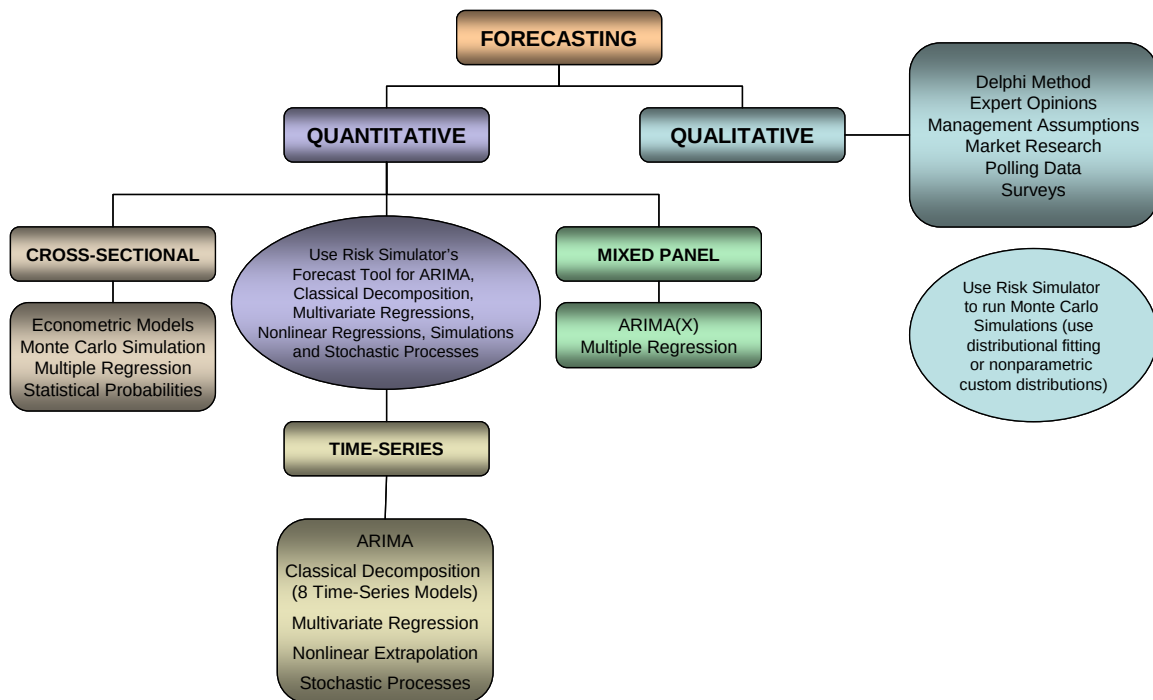


Figure 3.1 – 予測過程

様々なタイプの予測法の技術

予測法は大抵量的と質的に分けることができます。質的予測法は、確かでない、または存在しない履歴、同時期、および比較データが存在する時に使用されます。Delphi、および専門意見（もの領域の専門家、マーケティングの専門家、あるいは内部スタッフメンバーによる合意形成）、管理の仮定（上級管理者による目標成長率の設定）、マーケットリサーチ、外部データ、あるいは投票、および調査（第三者の情報源、産業やセクターの指標、および積極的市場調査などから得たデータ）のような複数の定性的予測法

があります。これらの推定は、単一・ポイントの推定（平均値の一致）か予測値の集合設定（予測の分布）のどちらかである可能性があります。その後、リスクシミュレーターにカスタム分布を記入することが出来、結果として生じる予測をシミュレーションすることが出来ます。これは、分布として推定されたこれらのデータポイント自体を使用したノンパラメトリックシミュレーションです。

定量的予測法では、実在するデータまたは予測を実行するのに必要なデータは、時系列（異なった年の収入、インフレーションの比率、利率、マーケットシェア、故障率等の時間の要素を持った値）、クロスセクション（各学生の SAT のレベル、IQ と一週間に消費するアルコール飲料の数が与えられ、ある特定の年の全国の大学2年生の平均の成績点のように時間から独立している値）、または混合パネル（時系列とパネル（補助的）データの混合、例えば、所与のマーケティングの費用とマーケットシェアの下での次期 10 年間の売上の予測、これは、売上が時系列であるけれども、マーケティングの費用やマーケットシェアが予測をモデル化する際の支援をするために存在することを意味する）に分割可能である。

データ

リスクシミュレーターのソフトウェアは様々な予測法を含んでいます。

1. ARIMA (自己回帰和分平均移動)
2. 自動 ARIMA
3. 自動計量法
4. 基本的な計量法
5. キュービック・スプライン
6. ニューラル・ネットワーク
7. GARCH (一般化自己回帰条件付き分散不均一モデル)
8. J 曲線
9. S 曲線
10. マルコフ連鎖
11. 組み合わせ的なファジィ論理
12. 最尤法
13. 非線形外押法
14. 回帰分析

- 15. 確率予測法
- 16. 時系列分析
- 17. トレンドライン

各予測法の分析の詳細は、このユーザーマニュアルの範囲から外れてしまいますが、詳しくは、リスクシミュレーターのソフトウェアの開発者であるジョナサン・マン博士によるリスクのモデル化：モンテカルロシミュレーション、リアルオプション分析、確率予測法とポートフォリオの最適化の適用(Wiley Finance 2006)を参照願います。それでも尚、下記に最も一般的なアプローチが表示されています。他のすべての予測法のアプローチはかなり簡単にリスクシミュレーターの中で適用する事ができます。

リスクシミュレーターで予測ツールを実行するにあたって

一般的に予測を作成するには複数のステップを行う必要があります。

- Excel を起動し記入するか、実在する履歴的データを開いてください。
- データを選択し、シミュレーション (Simulation) 、予測 (Forecasting) を選択してください。
- 適切な範囲 (ARIMA, 多変数回帰、非線形外押法、確率予測、時系列分析)を選択し、適切な入力を記入してください。

Figure 3.2 は、予測ツールと様々な方法を表示しています。



Figure 3.2 – リスク シミュレーターの予測過程

次に各方法の概略的検討とソフトウェアを使用するにあたって複数の簡潔な例証が与えられます。例証のファイルは、スタートメニューから開くことができます。[スタート \(Start\)](#) | [リアルオプションズバリュエーション \(Real Options Valuation\)](#) | [リスクシミュレーター\(Risk Simulator\)](#) | [例証\(Examples\)](#) を選択するか、[リスクシミュレーター\(Risk Simulator\)](#) | [例証モデル \(Example Models\)](#) を選択してください。

時系列分析

理論・セオリー:

Figure 3.3 は、季節性と傾向によって分離された八つの最も一般的な時系列モデルを表示しています。例えば、データ変数が傾向、および季節性を持っていない場合、単一移動平均モデル、または単一指数平滑法モデルで足りるはずです。但し、季節性が存在するが明らかな傾向が存在しない場合は、季節性の加法、または季節性の相乗的なモデルの適用がもっと適切だということです。

	No Seasonality	With Seasonality
No Trend	Single Moving Average	Seasonal Additive
	Single Exponential Smoothing	Seasonal Multiplicative
With Trend	Double Moving Average	Holt-Winter's Additive
	Double Exponential Smoothing	Holt-Winter's Multiplicative

Figure 3.3 – 八つの最も一般的な時系列法

手順:

- Excel を起動し、必要な場合は履歴的データを開いてください（下記の例証では例証フォルダーにある時系列予測のファイルを使用しています）。
- 履歴的データを選択してください（データは一行に表示されていなければいけません）

■ リスクシミュレーター(Risk Simulator) | 予測 (Forecasting) | 時系列分析(Time-Series Analysis)を選択してください。

■ 適用するモデルを選択し、関連する仮定を記入し、OK をクリックしてください。

履歴的な販売収入

年	4 部位数	周期	販売
2006	1	1	\$684.20
2006	2	2	\$584.10
2006	3	3	\$765.40
2006	4	4	\$892.30
2007	1	5	\$885.40
2007	2	6	\$677.00
2007	3	7	\$1,006.60
2007	4	8	\$1,122.10
2008	1	9	\$1,163.40
2008	2	10	\$993.20
2008	3	11	\$1,312.50
2008	4	12	\$1,545.30
2009	1	13	\$1,596.20
2009	2	14	\$1,260.40
2009	3	15	\$1,735.20
2009	4	16	\$2,029.70
2010	1	17	\$2,107.80
2010	2	18	\$1,650.30
2010	3	19	\$2,304.40
2010	4	20	\$2,639.40



Figure 3.4 – 時系列分析

結果の解釈にあたって:

Figure 3.5 は、予測ツールを通して生成されたサンプル結果を表示しています。使用されたモデルは、Holt-Winters の相乗的モデルです。 Figure 3.5 では、適合モデルと予測チャートは、傾向と季節性は Holt-Winters の相乗的モデルによってうまく選ばれている事を示しているのに注目してください。時系列分析のレポートは、重要な最適化されたアルファ、ベータとガンマのパラメーター、エラーの測定、適合したデータ、予測値と適合した予測のグラフを与えてくれます。パラメーターは単に基準として表示されています。アルファは、終止時間に変換する基準レベルの記憶の効果を取得し、ベータは、傾向の強さを測定する傾向パラメーターであり、ガンマは、履歴的データの季節性の強さを測定します。分析は、履歴的データをこれらの三つの要素に分解し、未来の予測を実行する為に構成しなします。適合されたデータは、構成のモデルを使用して適合したデータとして履歴的データを表示し、過去に対してどれくらい予測が近いのか（この技術はバックキャスティングと呼ばれています）を表示しています。 予測の値は、単一ポイ

ントの推定か仮定かの（シミュレーションプロファイルが存在し、仮定の自動生成が選択されている場合）どちらかです。グラフは、これらの履歴的、適合された予測の値を表示しています。チャートは強力なコミュニケーションであり、予測モデルがどれほど良いのかを見分ける視覚ツールでもあります。

メモ:

この時系列分析のモジュールには、Figure 3.3 で表示された八つの時系列モデルが含まれています。傾向と季節性の基準に基づいて特定のモデルを選択し実行することが出来、また、自動的に八つのモデルと反復的調整し、パラメーターを最適化し、データに最も適切なモデルを見出す自動モデルセクションを選択することも出来ます。一方、八つのモデルの中から一つのモデルを選択した場合、最適化のチェックボックスの選択を辞める事が出来、独自のアルファ、ベータ、そしてガンマパラメーターを入力することが可能となります。これらのパラメーターの技術的詳細には、ジョナサン・マン博士のリスクのモデル化：モンテカルロ・シミュレーションの適用、リアルオプションズ分析、予測と最適化 (Wiley, 2006) をご覧ください。また、自動モデルセクション、および季節性モデルのどれかを選択した場合、適切な季節性期間を入力する必要があります。季節性の入力には正の整数でなければいけません (例、データが 4 期の場合、年間の季節、サイクルの数に 4 と入力するか、データが月間の場合は 12 と入力してください)。次に予測する期間の数を記入してください。この値は、正の整数でなければいけません。最高の実行時間は 300 秒です。普通、変換の必要がありません。但し、かなりの量の履歴的データを基にした予測の分析が少しわずかに長くなり、過程時間が実行時間を超えた場合、過程は終止されてしまいます。また、自動的に過程の生成を行う予測を選択することが出来ます。これは、単一ポイントの推定の代わりに予測が仮定となります。最後に、極パラメーターオプションは、アルファ、ベータとガンマパラメーターを最適化する可能性を与え、0 と 1 を含む事が出来ます。一部の予測のソフトウェアしか、これらの極端なパラメーターの使用を可能にしてくれません。リスクシミュレーターは、どれを使用するか選択する可能性を与えてくれます。一般的には極端なパラメーターを使用する必要はありません。

Holt-Winterの相乗

統計の概略

アルファ, ベータ, ガンマ	RMSE	アルファ, ベータ, ガンマ	RMSE
0.00, 0.00, 0.00	914.824	0.00, 0.00, 0.00	914.824
0.10, 0.10, 0.10	415.322	0.10, 0.10, 0.10	415.322
0.20, 0.20, 0.20	187.202	0.20, 0.20, 0.20	187.202
0.30, 0.30, 0.30	118.795	0.30, 0.30, 0.30	118.795
0.40, 0.40, 0.40	101.794	0.40, 0.40, 0.40	101.794
0.50, 0.50, 0.50	102.143		

分析の実行に使用されたアルファ=0.2429, ベータ=1.0000, ガンマ=0.7797, と、季節性=4

時系列分析の概略

季節性と傾向が存在する時は、基本要素内でデータを分解する為により上級なモデルが必要となります；アルファパラメーターで量った基本ケースレベル(L)、ベータパラメーターで量った傾向コンボット(b)とガンマパラメーターで量った季節性コンボット(S)です。他の方法も存在しますが、最も一般的な二つの方法はHolt-Winterの相乗的季節性法とHolt-Wintersの相乗的季節性法です。Holt-Wintersの相乗的モデルでは基本ケースレベル、季節性と傾向は適合な予測を得る為と一緒に含まれています。

移動平均予測の最良の適合テストは、平均平方根誤差 (RMSE) を使用します。RMSEは、適合された値に対しての現在のデータポイントの平均平方偏差の平方根を計算します。

平均事情誤差 (MSE)は、エラーを二乗する絶対的なエラーの測定します (現在の履歴的データとモデルによって予測された予測適合データの間の違い)。これによって、各結果のキャンセルによって肯定的なエラーを維持します。またこの測定は、エラーを二乗することによって小さなエラーより大きなエラーを重大化することで大きなエラーを大げさに捉える傾向があります。これは、様々な時系列モデルを比較するのに手助けをしてくれます。平均平方二乗誤差 (RMSE) は、MSEの二乗平方で、もっとも典型的なエラーの測定です。ちなみに二次損失関数とも知られています。RMSEは、予測エラーの絶対値の平均として定義することが出来ます。また、予測エラーの費用と予測エラーの絶対サイズが比例している場合にもっとも適切です。RMSEは、時系列モデルの最適な適合の為に基準範囲として使用されています。

平均絶対百分位誤差 (MAPE)は、履歴的データポイントの平均百分位誤差として相対的な誤差統計測定で、予測エラーの費用がエラーの数値的サイズよりパーセントエラーが相対的に密接な時により適切です。最後に、TheilのU統計は関連測定で、モデルの予測の単純な思考を測定します。これは、TheilのU統計が1.0未満の場合、使用された予測方法は統計的に推定を行い、当てだすよりも確実なことを示します。

期数	現在	適合する予測	エラーの測定
1	684.20		RMSE 71.8132
2	584.10		MSE 5157.1348
3	765.40		MAD 53.4071
4	892.30		MAPE 4.50%
5	885.40	684.20	TheilのU 0.3054
6	677.00	667.55	
7	1006.60	935.45	
8	1122.10	1198.09	
9	1163.40	1112.48	
10	993.20	887.95	
11	1312.50	1348.38	
12	1545.30	1546.53	
13	1596.20	1572.44	
14	1260.40	1299.20	
15	1735.20	1704.77	
16	2029.70	1976.23	
17	2107.80	2026.01	
18	1650.30	1637.28	
19	2304.40	2245.93	
20	2639.40	2643.09	
予測21		2713.69	
予測22		2114.79	
予測23		2900.42	
予測24		3293.81	

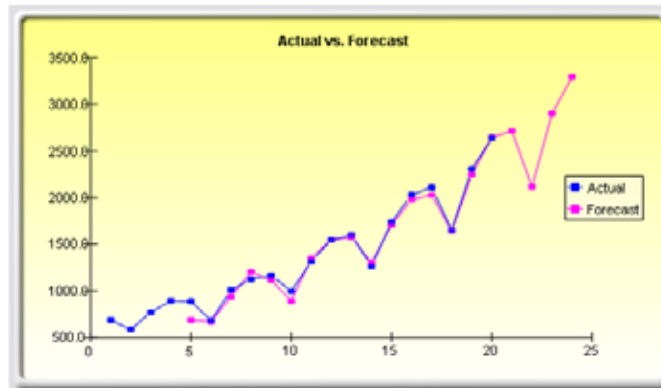


Figure 3.5 – 例証：Holt-Wintersの予測レポート

理論・セオリー:

ここでは、ユーザーが回帰分析の基礎について十分な知識を持っていることが必要となります。一般的な 2 変数的な線形回帰の方式は $Y = \beta_0 + \beta_1 X + \varepsilon$ で、 β_0 は切片であり、 β_1 は傾きで、 ε はエラーを示しています。これは、二つの変数のみを含む 2 変数関数で、 Y または従属変数と、 X または独立変数となります。 X は、回帰変数とも知られています (時折、2 変数回帰は、独立変数 X だけが存在する時は単回帰と呼びます)。従属変数は、独立変数に依存する為従属変数と呼ばれています。例えば、販売の収入は、商品の広告やプロモーションに使用されたマーケティングの費用に従属しており、従属変数の販売と独立変数のマーケティングの費用を作成します。2 変数回帰の例証は、Figure 3.6 の左のパネルに表示されている 2 次元面でのデータポイントの設定を通して、最適な適合線の挿入です。他のケースでは、複数、または独立した X 変数の n 数の多変数回帰を実行できます。普通の回帰方式は $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 \dots + \beta_n X_n + \varepsilon$ です。このケースでは、最適な適合線は、 $n + 1$ 次元面内に含まれます。

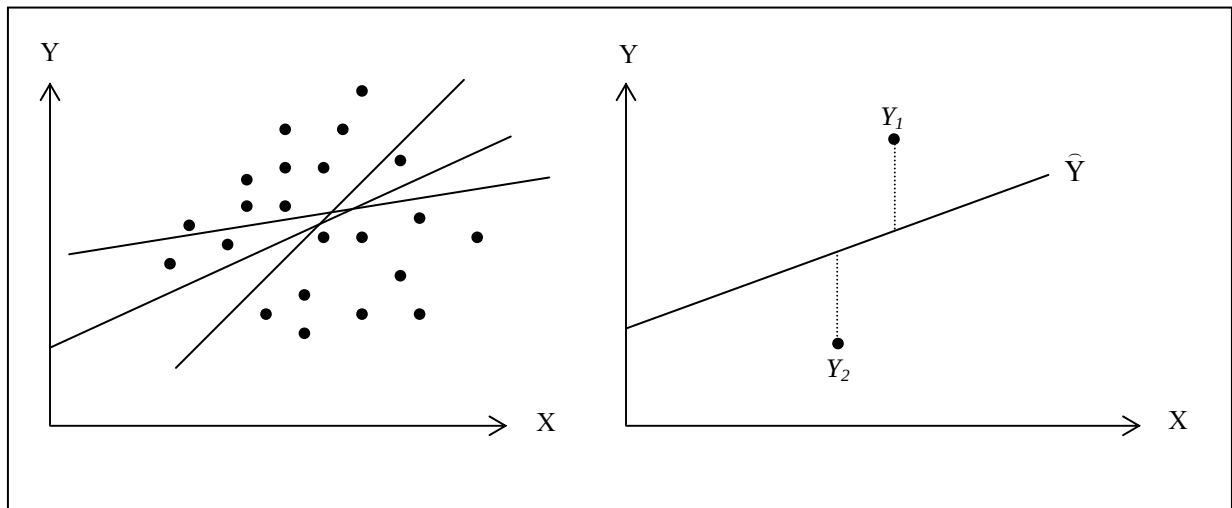


Figure 3.6 – 二変量回帰

但し、Figure 3.6 の様にデータポイントの設定を通して線を適合すると、結果として複数の確率的な線が表示されます。最適な適合の線は、総合的な縦軸のエラーを最小化する単一的な特異な線として定められます。これは、Figure 3.6 の右側のパネルで表示さ

れているように、現在のデータポイント(Y_i)と、予期された線 ($\hat{Y}$)の間の絶対距離の足し算です。エラーを最小化する最適な適合の線を見出す為には、もっと上級のなアプローチが必要となります。これが回帰分析です。よって回帰分析は、全体的なエラーを最小化を要する次式のような計算によって特異な最適の適合線を見出します。

$$\text{Min} \sum_{i=1}^n (Y_i - \hat{Y}_i)^2$$

ここでは、一つの特異な線だけが偏差平方を最小化します。エラー（現在のデータポイントと予期された線の間の縦軸の距離）は、正のエラーを負のエラーがキャンセルする事を避ける為に平方処理されます。傾きと切片に関する最小化の問題の解決には、1次導関数を計算し、ゼロに等しくなるように計算をしなければいけません:

$$\frac{d}{d\beta_0} \sum_{i=1}^n (Y_i - \hat{Y}_i)^2 = 0 \text{ and } \frac{d}{d\beta_1} \sum_{i=1}^n (Y_i - \hat{Y}_i)^2 = 0$$

この結果は、2変数回帰の最小二乗公式をもたらします:

$$\beta_1 = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{\sum_{i=1}^n (X_i - \bar{X})^2} = \frac{\sum_{i=1}^n X_i Y_i - \frac{\sum_{i=1}^n X_i \sum_{i=1}^n Y_i}{n}}{\sum_{i=1}^n X_i^2 - \frac{\left(\sum_{i=1}^n X_i\right)^2}{n}}$$

$$\beta_0 = \bar{Y} - \beta_1 \bar{X}$$

多変数回帰にも、複数の独立変数を考慮して類似の拡張を行うことができ、すなわち、

$Y_i = \beta_1 + \beta_2 X_{2,i} + \beta_3 X_{3,i} + \varepsilon_i$ で、その結果、推定傾斜は、は次の通りに計算されます:

$$\hat{\beta}_2 = \frac{\sum Y_i X_{2,i} \sum X_{3,i}^2 - \sum Y_i X_{3,i} \sum X_{2,i} X_{3,i}}{\sum X_{2,i}^2 \sum X_{3,i}^2 - \left(\sum X_{2,i} X_{3,i}\right)^2}$$

$$\hat{\beta}_3 = \frac{\sum Y_i X_{3,i} \sum X_{2,i}^2 - \sum Y_i X_{2,i} \sum X_{2,i} X_{3,i}}{\sum X_{2,i}^2 \sum X_{3,i}^2 - \left(\sum X_{2,i} X_{3,i}\right)^2}$$

多変数回帰を実行するにあたって、結果の設定と解釈には、注意が必要です。例えば、計量経済学的なモデル化の際には優れた理解が (例、構造の破綻、多重共線性、不均一分散、自己相関、特定のテスト、非線形などのような落とし穴の検出)適切なモデルの

構築には必要となります。多変数回帰での分析と議論及びどのようにしてこれらの落とし穴を検出するかの詳細には、ジョナサン・マン博士によるリスクのモデル化: モンテカルロ・シミュレーションの適用、リアルオプションズ分析、予測と最適化 (Wiley, 2006) を参照ください。

手順:

- Excel を起動し、必要な場合は履歴的データを開いてください (下記の表では、例証ファイルの**複数の回帰**を使用しています)。
- 編集されたデータが列にある事を確認してください。変数名称を含めて全てのデータ範囲を選択し、**リスクシミュレーター(Risk Simulator) | 予測(Forecasting) | 複数の回帰 (Multiple Regression)**を選択してください。
- 従属変数を選択し、重要なオプション (タイムラグ、ステップワイズ回帰、非線形回帰等)を選択し、OK をクリックしてください。

結果の解釈:

Figure 3.8 は、サンプルの多変数回帰結果のレポートを表示しています。レポートは全ての回帰結果、分散分析結果、適合チャートと仮説検定の結果が含まれた完全な情報を表示してくれます。これらの結果の解釈の技法の詳細は、このユーザーマニュアルの最後にあります。回帰レポートの解釈のように多変数回帰の議論と分析の詳細には、ジョナサン・マン博士のリスクのモデル化: モンテカルロ・シミュレーションの適用、リアルオプションズ分析、予測と最適化(Wiley 2006) を参照ください。

重回帰

Y	X1	X2	X3	X4	X5
521	18308	185	4.041	79.6	7.2
367	1148	600	0.55	1	8.5
443	18068	372	3.665	32.3	5.7
365	7729	142	2.351	45.1	7.3
614	100484	432	29.76	190.8	7.5
385	16728	290	3.294	31.8	5
286	14630	346	3.287	678.4	6.7
397	4008	328	0.666	340.8	6.2
764	38927	354	12.938	239.6	7.3
427	22322	266	6.478	111.9	5
153	3711	320	1.108	172.5	2.8
231	3136	197	1.007	12.2	6.1
524	50508	266	11.431	205.6	7.1
328	28886	173	5.544	154.6	5.9
240	16996	190	2.777	49.7	4.6
286	13035	239	2.478	30.3	4.4
285	12973	190	3.685	92.8	7.4
569	16309	241	4.22	96.9	7.1
96	5227	189	1.228	39.8	7.5
498	19235	358	4.781	489.2	5.9
481	44487	315	6.016	767.6	9
468	44213	303	9.295	163.6	9.2
177	23619	228	4.375	55	5.1
198	9106	134	2.573	54.9	8.6
458	24917	189	5.117	74.3	6.6
108	3872	196	0.799	5.5	6.9
246	8945	183	1.578	20.5	2.7
291	2373	417	1.202	10.9	5.5
68	7128	233	1.109	123.7	7.2
311	23624	349	7.73	1042	6.6
606	5242	284	1.515	12.5	6.9
512	92629	499	17.99	381	7.2
426	28795	231	6.629	136.1	5.8
47	4487	143	0.639	9.3	4.1
265	48799	249	10.847	264.9	6.4
370	14067	195	3.146	45.8	6.7



1. ヘッダーを含めたデータ範囲を選択してください。(B5:G55)
 2. リスクシミュレータ (Risk Simulator) | 予測 (Forecasting) | 重回帰 (Multiple Regression)
 3. 従属変数を選択し (この例では、Y変数)、必要に応じて特定の变换を行ってください (時差回帰、非線形回帰、ステップワイズ回帰)
- そして、OKをクリックしてください。分析結果については回帰レポートをご覧ください。



Figure 3.7 – 多変量回帰の実行

回帰分析レポート

回帰統計

R-Squared (決定係数)	0.3272
調整された決定係数	0.2508
複数の R (複数の相関係数)	0.5720
測定での標準エラー (SEy)	149.6720
観測	50

決定係数(RSquared)、または、定義の相関は、従属変数の変化の0.33は、この回帰分析で独立変数によって計算、そして帰説明する事ができることを示しています。しかし、重回帰では、調整された決定係数(RSquared)は、実存する追加の独立変数の計算に含まれ、また、回帰と決定係数をより正確で納得のいく値に調整します。したがって、従属変数の0.25の変化だけが回帰によって説明できます。

重相関係数(Multiple R)は、回帰方程式に基づいた、現在の従属変数(Y)推定、または適合(Y)の間の相関関係を測定します。これは、決定の相関(R-Square)の二乗根です。

推定の標準エラー(SEy)は、回帰線、または平面を下回る、および、上回るデータポイントの離散を記述します。この値は、推定後の信頼区間を得るための計算の一部に使われます。

回帰結果

	インタセプト	学士号相当	一人当たりの 警察の費用	人口(数百万)	人口密度(人 口/平方マイル)	失業率
相関	57.9555	-0.0035	0.4644	25.2377	-0.0086	16.5579
標準エラー	108.7901	0.0035	0.2535	14.1172	0.1016	14.7996
t-統計	0.5327	-1.0066	1.8316	1.7877	-0.0843	1.1188
p-値	0.5969	0.3197	0.0738	0.0807	0.9332	0.2693
5%より下	-161.2966	-0.0106	-0.0466	-3.2137	-0.2132	-13.2687
95%より上	277.2076	0.0036	0.9753	53.6891	0.1961	46.3845

自由度

	自由度	仮説
回帰の為の自由度	5	批判的な t-統計 (99% の信頼度と 44の自由度)
残差のための自由度	44	批判的な t-統計 (95% の信頼度と 44の自由度)
合計自由度	49	批判的な t-統計 (90% の信頼度と 44の自由度)

相関は、推定された回帰のインタセプトと斜面を与えてくれます。例えば、相関が本当の見積もりの場合、母集団 bの値は次の方程式に当てはまります。Y = b0 + b1X1 + b2X2 + ... + bnXn。標準エラーはどれほど確率的に相関を予測したかを推定します。また、t-統計は、これらの標準エラーを基に予測された各相関の比率です。

t-統計は、本当の相関の平均値 = 0の帰無仮説(Ho)、本当の平均値が0でない、対立仮説(Ha)が設定された、仮説検定に使われます。T-検定が実行され、計算されたt-統計は、批判的な値として関連した残差の自由度と比較されます。他の回帰を対象に、各相関が統計的に有意だった場合、t-検定そのものと、t-検定の計算はとっても重要なものとなります。これは、t-検定は回帰関数、または独立変数が回帰に残されるべきか、またはドロップされるべきかを統計的に確認します。

相関は、計算された批判的なt-統計の有意な自由度(df)に対して統計的に上回る場合は統計的に有意な事を示します。三つの主な信頼レベルは、90%、95%と99%の有意度をテストする為に使用されます。t-統計が信頼レベルを上回った場合、統計的に有意だと解釈されます。また、p-値の計算は、小さいp-値ほど、相関が有意であるなどの、各t-統計の発生の確立を表示します。普段のp-値の有意レベルは、0.01、0.05と0.10に連応した99%、95%と99%の信頼レベルです。

青のハイライトで表示されたp値と相関は、これらが90%の信頼レベル、または0.10のアルファレベル、統計的に有意である事を示しています。一方、赤のハイライトで表示された範囲はどのアルファレベルに対しても統計的に有意でない事を示しています。

変化の分析

	処理平方和	平方の平均値	F-統計	p-値	仮説検定	
回帰	479398.4898	95877.6990	4.2799	0.0029	拒否的なF-統計 (90% の信頼度、4 と 30の自由度)	3.4651
残差	985675.1902	22401.7089			拒否的なF-統計 (90% の信頼度、4 と 30の自由度)	2.4270
合計	1465063.6800				拒否的なF-統計 (90% の信頼度、4 と 30の自由度)	1.9828

変化表(ANOVA)の分析は、総合統計的有効な回帰モデルのF-検定を表示してくれます。F検定のように、一つずつの回帰を計算する代わりに、F検定は測定した全ての有価の統計的特徴を見出してくれます。F統計は、平方の回帰の平均値と平方の残差の平均の比率として計算されます。分子は、どれほど回帰が説明されているかを測定している間、位分母はどれほど説明されていないかを測定します。したがって、大きなF統計は、モデルがより有意という事です。適切なp値は、全ての相関がゼロと同一している帰無仮説(H0)をテストする為に計算されます。一方、対立仮説(Ha)は、全ての値がゼロと異なっていて、全体的に回帰モデルが有意なことを示しています。これによって、p値がアルファ有意の0.01、0.05、または0.10より小さい場合は、回帰が有意だという事を示します。様々な有意レベルで計算されたF-統計と批判的なFの値の比較によって同様な適応がF統計に当てはまります。

予測

Period	Actual (Y)	Forecast (F)	Error (E)
1	521	299.5124	221.4876
2	367	487.1243	(120.1243)
3	443	353.2789	89.7211
4	365	276.3296	88.6704
5	614	778.1336	(162.1336)
6	305	290.9993	86.0007
7	286	354.8718	(68.8718)
8	397	312.6155	84.3845
9	764	529.7550	234.2450
10	427	347.7034	79.2966
11	153	266.2526	(113.2526)
12	231	264.6375	(33.6375)
13	524	406.8009	117.1991
14	328	272.2228	55.7774
15	240	231.7082	8.2118
16	286	257.8862	28.1138
17	205	314.9521	(29.9521)
18	589	335.3140	233.6860
19	96	282.0356	(186.0356)
20	498	370.2062	127.7938
21	401	340.0742	140.1258
22	468	427.5118	40.4882
23	177	274.5290	(97.5290)
24	198	294.7795	(96.7795)
25	458	295.2180	162.7820
26	108	269.6195	(161.6195)
27	248	195.5955	50.4045
28	291	364.5004	(73.5004)
29	68	287.0428	(219.0428)
30	311	431.7568	(120.7568)
31	608	323.6389	282.3601
32	512	531.4356	(19.4356)
33	426	325.3641	100.6359
34	47	192.3960	(145.3960)
35	265	378.1250	(113.1250)
36	370	288.6064	81.3936
37	312	317.5374	(5.5374)
38	222	355.0075	(133.0075)
39	280	316.6280	(36.6280)
40	759	301.1523	457.8477

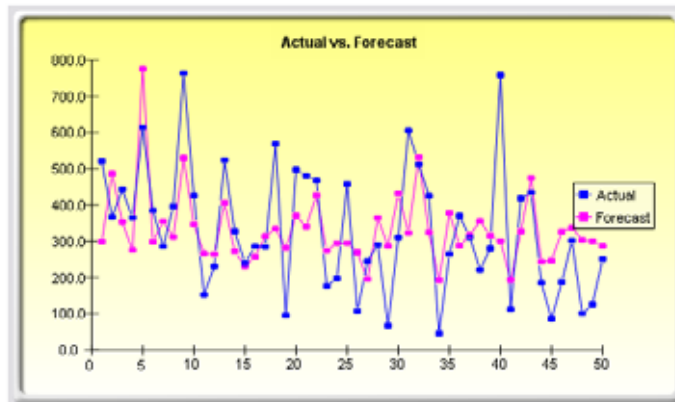


Figure 3.8 – 多変量回帰の結果

確率的予測法

理論・セオリー:

確率過程は、時間が経つに連れて作成される結果のシリーズの数学的な公式でしかありません。この結果は性質として決定論ではありません。これは、毎年上昇する X パーセントの価格やこの X の要素に足される Y パーセントによって上昇する収入のように明らかでシンプルなルールに従う方式、または過程ではありません。確率過程とは、非

決定論である事から呼ばれ、確率過程の公式内に数値を繋ぐ事が出来、毎回異なった結果を得ることが出来ます。例えば、株式価格のルートは性質として確率的である為、確実な株式価格のルートを予測することが出来ません。但し、時間が経つに連れて進化する価格は、これらの価格を生成する過程で囲まれます。過程は結果と違って、事前に固定され、定められています。したがって、確率シミュレーションによって、価格の複数の経路を作成し、これらのシミュレーションの統計的なサンプリングを取得し、時系列の生成に使用された確率過程のパラメーターと性質を所与とした場合に実際の価格が取り得る潜在的な経路に関して推測することが出来ます。三つの基本的な確率過程がリスクシミュレーターの予測ツールに含まれており、先ず、幾何ブラウン運動あるいはランダムウォークは、単純で幅広い応用によって最も一般的で広く使用されています。他の二つの確率過程は、平均回帰過程とジャンプ拡散過程です。

確率過程シミュレーションでの興味深いところは、履歴的データが必ずしも必要ないということです。これは、モデルは履歴データの集合に適合する必要があると言う事です。ただ、履歴データのボラティリティと期待リターンを計算するだけか、外部データと比較してそれらを推定するか、これらの値の仮定を作成するだけと言う事です。どのようにして各入力計算（例、平均回帰の比率、ジャンプ確率、ボラティリティ等）されるのか等の詳細には、ジョナサン・マン博士のリスクのモデル化: モンテカルロ・シミュレーションの適用、リアルオプションズ分析、予測と最適化、第2版 (Wiley 2006)を参照してください。

手順:

- **リスクシミュレーター(Risk Simulator) | 予測 (Forecasting) | 確率過程 (Stochastic Processes)**を選択することでモジュールを起動してください。
- 希望する過程を選び、必要とする入力を記入してください。数回チャートを更新するにクリックし、過程が希望通りの振る舞いを表示していることを確認し、OK をクリックしてください。(Figure 3.9)

結果の解釈:

Figure 3.10 は、サンプルの確率過程を表示しています。チャートは、反復のサンプル設定を表示しており、レポートは、確率過程の基礎を記述しています。また、各時間の期間の為の予測値（平均値と標準偏差）が与えられています。これらの値の使用によって、分析にとってどの時間周期が一番重要か定めることが出来、正規分布を使用したこれらの平均と標準偏差の値に基づいた仮定の設定が出来ます。これらの仮定は後ほど、カスタムモデルでシミュレーションすることが出来ます。

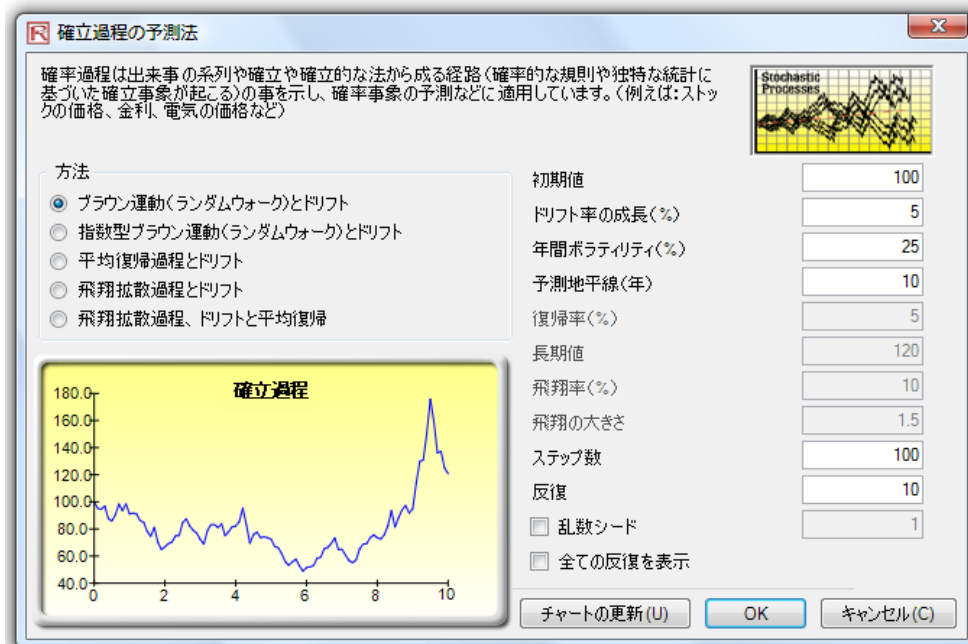


Figure 3.9 – 確率過程方の予測

確率過程予測法

統計の概略

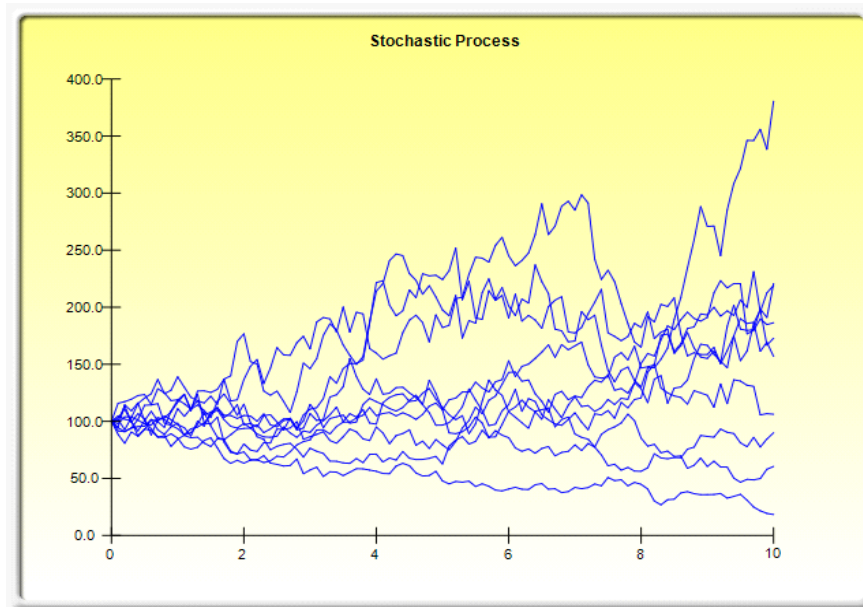
確率過程は、確率法から生成された連続事象、または連続パッドの事を言います。これは、ランダムな事象は長期間に発生する事がありますが、特定の統計的な、そして確率的なルールによって定められています。主な確率過程は、ランダムウォーク、またはブラウン運動、平均回帰とJumpDiffusionが含まれています。これらの過程は、見た目ランダムな傾向を追うが、確立法に制限されている複数の変数を予測する事が出来ます。リスク・シミュレーターの確率仮定モジュールで各過程を作成、シミュレーションできます。

ランダムウォーク、ブラウン運動過程は、株価、物価などの確立時系列データが与えられたドリフト、または成長率とドリフトパッド周辺のボラティリティなどの予測に適切しています。平均回帰過程は、パッドで長期値を目標に切り替える事で、ランダムウォーク仮定の変動の減少の為に使われます。利率やインフレ率などの長期値を持った時系列の変数の予測をしやすいからです(これらは、マーケットか、監督機関で長期目標率です)。JumpDiffusion過程は、変数が時折、ランダムな飛翔の可能性を持っている時の時系列データを予測する為に適切しています。例えば、石油価格や電気価格です(離散的な外因性のイベントショックは、価格を飛翔させたり、落とさせたりします)。最後に、これらの三つの確率過程は要求に応じて適合させたり、混合したり出来ます。

右側に表示されている結果は、全てのタイム・ステップの反復で生成された平均値と標準偏差を示しています。もし、”全ての反復表示”がオプションで選択されている場合は、各反復が別々のワークシートに表示されます。下に生成されるグラフは、反復のサンプル設定が表示されます。

確率過程: ブラウン運動 (ランダムウォーク) とドリフト

開始地	100	ステップ	100.00	飛翔率	N/A
ドリフト率	5.00%	反復	10.00	ジャンプサイズ	N/A
ボラティリティ	25.00%	復帰率	N/A	乱数シード	365705039
地平線	10	長期間の値	N/A		



時間	平均値	標準偏差
0.0000	100.00	0.00
0.1000	99.25	7.79
0.2000	101.72	12.40
0.3000	101.39	8.28
0.4000	102.46	12.60
0.5000	104.08	11.16
0.6000	103.52	11.71
0.7000	105.05	17.00
0.8000	101.33	13.95
0.9000	102.07	14.21
1.0000	106.31	17.48
1.1000	103.28	18.63
1.2000	99.48	14.86
1.3000	105.66	16.82
1.4000	102.26	14.03
1.5000	102.93	14.63
1.6000	104.56	16.59
1.7000	102.49	24.04
1.8000	97.96	23.85
1.9000	100.67	30.51
2.0000	103.94	33.05
2.1000	99.74	29.96
2.2000	97.79	31.16
2.3000	94.69	22.79
2.4000	93.62	26.00
2.5000	97.34	29.60
2.6000	96.83	27.18
2.7000	97.18	25.95
2.8000	98.22	29.49
2.9000	100.66	36.46
3.0000	103.63	32.15
3.1000	107.32	37.65
3.2000	108.97	43.25
3.3000	113.68	46.07
3.4000	112.80	42.53
3.5000	115.06	46.09
3.6000	115.80	42.18
3.7000	115.80	42.32
3.8000	116.97	42.81
3.9000	120.26	45.96
4.0000	127.50	56.05
4.1000	126.82	57.79
4.2000	126.73	58.84
4.3000	127.98	57.91
4.4000	131.95	57.87
4.5000	132.04	59.18
4.6000	128.87	59.84
4.7000	128.88	60.62
4.8000	129.78	60.18
4.9000	129.68	60.15
5.0000	122.27	59.31
5.1000	122.03	60.37
5.2000	129.89	68.13
5.3000	125.00	53.98

Figure 3.10 – 確率予測法の結果

非線形外押法

理論・セオリー:

外押法は、未来の特定の期間に予測された履歴的傾向の使用によって統計的な予測を作成します。これは、時系列だけで使用されます。クロスセクション、または混合パネルのデータ (クロスセクションデータを活用した時系列) にとっては、多変数回帰の方が適切です。この方法は、大きな変化が期待されない場合に最適です。これは、原因要素がと一定に維持すると期待される、または状況の原因要素が明確には知られていない時に

使用するのが適切です。また、過程への個人的なバイアスへの導入を防ぐのに役立ちます。外押法はかなり信頼でき、比較的シンプルで、経済的です。但し、近況、そして履歴的傾向が継続すると仮定する外押法は、プロジェクトの期間内で不連続が発生した場合、大きい予測エラーを生み出します。すなわち、時系列の純粋な外押法は、予測されるシリーズの履歴値に全ての必要な情報が含まれていると仮定しています。過去の振る舞いが未来の振る舞いを予言するのに最適だと仮定するならば、外押法が最適です。よって、全ての必要なことものは、多くの短い予測だという時に、とっても使いやすい、適切なアプローチです。

この技法は、全ての任意の x 値を通る滑らかな非線形曲線を挿入し、履歴的データ集合に未来の x 値を外押することで $f(x)$ の関数を推定します。この技法は、多項式の関数形態、またはロジスティック公式の形態（二つの多項式の比率）のどちらかを使用します。一般的に、データの良い振る舞いを得る為には、多項式の関数形態で十分ですが、ロジスティック関数形態は、時によってもっと精度を見せることがあります（特に極関数の時。例、分母がゼロに接近する関数）。

手順:

- Excel を起動し、必要な場合は履歴的データを開いてください (次に表示されているイラストは例証フォルダーの **非線形外押法** のファイルを使用しています)。
- 時系列データを選択し、 **リスクシミュレーター (Risk Simulator) | 予測 (Forecasting) | 非線形外押法 (Nonlinear Extrapolation)** を選択してください。
- 外押法のタイプを選択し(自動選択、多項式関数、および有理関数)、希望する予測期間の数を記入し (Figure 3.11)、OK をクリックしてください。

結果の解釈:

Figure 3.12 で表示されている結果は、外押された予測の値、エラーの測定法と外押法の結果のグラフを表示しています。エラーの測定法は、予測の妥当性を確認する為に使用出来、予測の性質を比較する時に最も重要となり、時系列分析に対して外押法の精度を比較する時に適切です。

メモ:

履歴的データが滑らかで幾つかの非線形のパターンと曲線に従っている時は、外押法は時系列分析より最適です。但し、データパターンが季節性のサイクルと傾向を辿っている時は時系列分析の方が良い結果を示すでしょう。

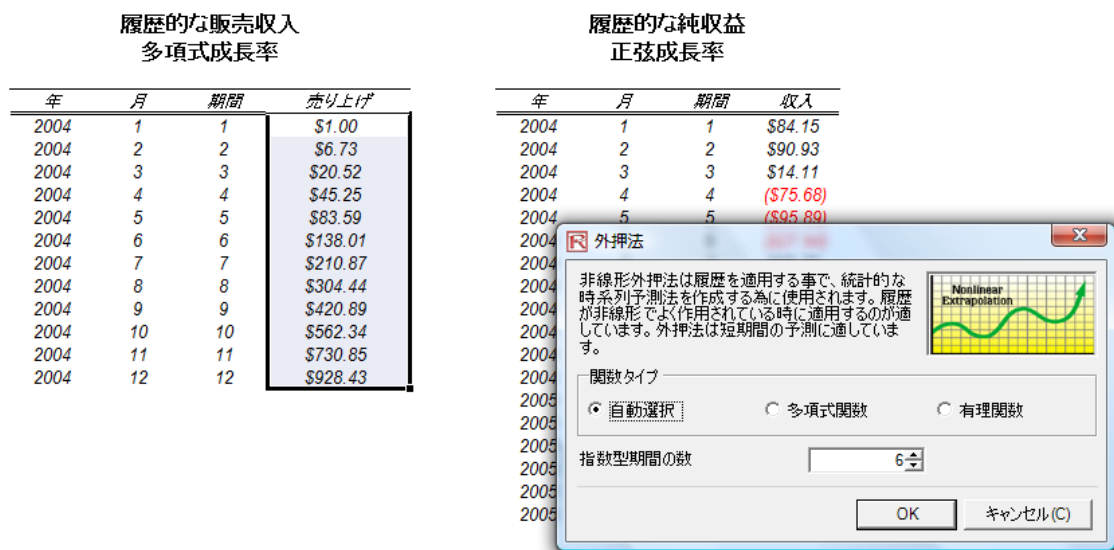


Figure 3.11 – 非線的外押法の実行

非線形外挿法

統計の概略

外挿法は、未来の特定の期間の状況を履歴の傾向を使用することで、統計的な予想を含みます。これは、時系列予測法でしか使用されません。横断、および混合パネルデータ（時系列と横断データ）には、重回帰が最適です。この方法は大きな変化が予想されていないときに適切です。これは、原因要因は残された定数に予測されているか、原因要因の状態がはっきり理解されない時などを示します。それはまたプロセス内での個人的な先入観への導入の失望等を援助します。外挿法はかなり信頼でき、比較的使やすく、安価です。ただし、最近、及び履歴の傾向は連続たと仮定する外挿法は、反映された期間内で不連続が発生すると、長期の予測エラーが表示されます。つまり、時系列の純粋な外挿法は、知る必要があるすべてのデータは、予測されたシリーズの履歴的価値に含まれていると仮定されています。過去の履歴的な振る舞いは未来の振る舞いの有望な予言だと仮定した場合、外挿法は大変な権威です。これは、全てのデータを短期間で予測が必要ときに適切です。

この方法は、すべてのxの値による滑らかな非線形カーブの挿入によってあらゆる任意のy値のためのf(x)関数を推定し、この滑らかなカーブを使用して、歴史的データセットを越える未来のxの値を外挿法で推定します。方法は多項式関数形式か理性的関数形式(2つの多項式の比率)を用います。普通、多項式関数形式は、よい振る舞いを持ったデータのために十分であるが、理性的関数形式は時によってもっと正確なときがあります(特に極関数と、例えば、ゼロに近づく分母と関数など)。

期間	現在	適合な予測	エラーの推定	エラーの測定
1	1.00			RMSE 19.6799
2	6.73	1.00		MSE 387.2974
3	20.52	-1.42	-8.15	MAD 10.2095
4	45.25	99.82	119.36	MAPE 31.56%
5	83.59	55.92	-46.67	Theilの U 1.1210
6	138.01	136.71	14.39	
7	210.87	211.96	1.69	
8	304.44	304.43	-0.41	
9	420.89	420.89	0.01	
10	562.34	562.34	0.00	
11	730.85	730.85	0.00	
12	928.43	928.43	0.00	
予測 13		1157.03	0.00	
予測 14		1418.57	0.00	
予測 15		1714.95	0.00	
予測 16		2048.00	0.00	
予測 17		2419.55	0.00	
予測 18		2831.39	0.00	

機能タイプ: 理性的

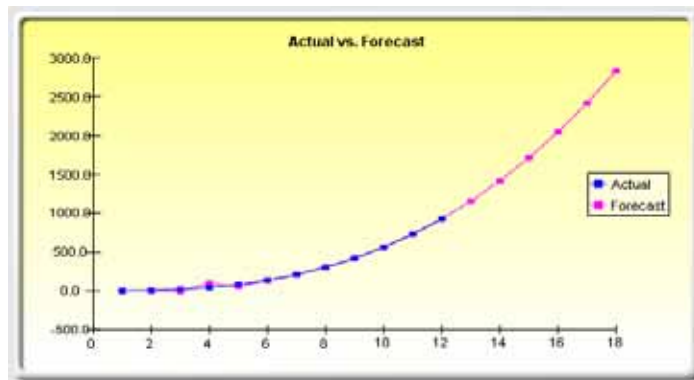


Figure 3.12 – 非線的外挿法の結果

Box-Jenkins ARIMA 高度な時系列

理論・セオリー:

一つのとっても高度で強力な時系列予測のツールは ARIMA、または自己回帰和分移動平均です。ARIMA の予測は三つの離れたツールを分かりやすいモデルの中に集めます。最初のツールセグメントは、自己回帰、および “AR” で、条件の無い予測モデルに残るタイムラグの値の数に対応します。本質的には、モデルは予測モデルへの現在のデータの履歴的変動を採取し、この変動あるいは残差を使ってもっと優れた予測モデルを作り出します。二つ目のツールセグメントは和分で “I” で表示されています。この和分は、予測される時系列を微分する回数に当たります。この要素は、データに存在する全ての非線形の成長の比率をカウントします。三つ目のツールセグメントは、移動平均、および “MA” で表示されており、本質的に時間差を持った予測エラーの移動平均を示してい

ます。この時間差を持った予測のエラーを含むことでモデルは、本質的にこれらの予測エラー、または間違いから学び、移動平均の計算を通して修正をします。ARIMA モデルは、Box-Jenkins 方法に従い、モデル構成から得た各ステップをランダムノイだけが残るまで行います。因みに、ARIMA のモデル化は、予測の生成に相関を使用します。ARIMA は、プロットされたデータで表示されないモデルパターンも使用できます。また、ARIMA モデルは、外生変数と混合する事が出来ませんが、外生変数が、追加の予測期間を予測するのに十分なデータポイントを持っている事を確認してください。最後に、モデルの複雑さの為、このモジュールは長期実行である事を理解して置いてください。

ARIMA モデルが一般的な時系列分析と多変数回帰より高度な理由が複数あります。一般的な時系列分析と多変数回帰の理解は、残ったエラーはこれら自身の時間差を持った値との相関です。この連続的な相関は、ディスタースは他のディスタースと相関しないという回帰理論の標準的な仮定に違反します。連続相関に関連した主要な問題には次の点が考えられます。

- 回帰分析と基本的な時系列分析は異なった線形の推定値の間ではもはやあまり有効ではありません。但し、エラーの残差が現在のエラーの残差を予測するのに役立つので、ARIMA を使用して従属変数のより良い予測にこの情報を活用することが可能となります。
- 回帰と時系列方式を使用して計算された標準エラーは正しくなく、一般的に控えめに述べられ、もし、回帰予測値として時間差のある従属変数が存在する場合、回帰の推定にはバイアスがあり、首尾一貫していませんが、ARIMA を使用する事によって修正する事が出来ます。

自己回帰和分移動平均、および ARIMA(p,d,q) モデルは、時系列データで連続相関のモデル化の為に三つの要素を使用する AR モデルの延長です。最初のコンポーネントは自己回帰(AR)です。AR(p)モデルは、公式での時系列の p の時間差を使用します。AR(p)モデルの方式は $y_t = a_1 y_{t-1} + \dots + a_p y_{t-p} + e_t$ です。二つ目のコンポーネントは和分 (d)です。各和分 $j = 1, 0$ いう $y = 7 - 9$; おは、時系列を微分するのに対応します。I(1)は、データを一度微分することを意味しています。I (d)は、d 回のデータの微分を意味しています。三つ目のコンポーネントは移動平均 (MA)です。MA(q) モデルは、予測を向上する為に予

測エラーの時間差 q を使用します。MA(q)モデルの方式は $y_t = e_t + b_1 e_{t-1} + \dots + b_q e_{t-q}$ です。最後に、ARMA(p, q) モデルのは結合形態 $y_t = a_1 y_{t-1} + \dots + a_p y_{t-p} + e_t + b_1 e_{t-1} + \dots + b_q e_{t-q}$ を持っています。

手順:

- Excel を起動し、データを記入するか実在するワークシートと履歴的データを予測するのに開いてください (イラストでは、次の例証で使用するファイル **時系列 ARIMA** を表示しています)。
- 時系列データを選択し、**リスクシミュレーター (Risk Simulator) | 予測 (Forecasting) | ARIMA** を選択してください。
- 重要な $P, D,$ と Q のパラメーター (正の整数のみ) を入力し、希望する予測周期の数を記入し、OK をクリックしてください。

結果の解釈:

ARIMA モデルの結果の解釈には、多変数の回帰分析のほとんどの指定と同じです (ARIMA モデルと多変数の回帰分析の解釈技法の詳細はジョナサン・マン博士のリスクのモデル化、第2版を参照ください)。但し、Figure 3.14 で表示されているように ARIMA 分析への特定の結果の様々な付加設定があります。初めに、ARIMA モデルの選択、および同一証明でよく使用される赤池情報基準 (AIC) とシュワルツ 基準 (SC) の付加です。これは、AIC と SC は、特定の p, d と q のパラメーターを持った特定のモデルが最適な適合の統計かどうかを定める為に使用されます。SC は、AIC より付加係数の為の大きなペナルティを与えますが、一般的には、AIC と SC の値が低いモデルが選択されます。最後に、自己相関 (AC) と呼ばれる結果の付加設定と部分自己相関 (PAC) 統計は、ARIMA レポートで与えられています。

例えば、自己相関 AC(1)がゼロでない場合、シリーズは一次連続的相関を意味します。もし、AC が幾何的な時間差の増加にて衰退する場合、シリーズは低次の自己回帰過程に従っている事を示しています。もし、少数の時間差の後に AC がゼロに下降した場合は、シリーズは低次の移動平均過程に従っている事を示しています。一方、PAC は、介入する時間差から相関を取り除いた後の k 周期離れた値の相関の測定をします。相関の

パターンが k より少ない次数の自己回帰によって採取された場合は、時間差 k における部分相関はゼロに近似します。Ljung-Box Q-統計と時間差 k に対してのこれらの p -値は既に与えられており、検定された帰無仮説は、 k の次数まで自己相関が見られません。自己相関のプロットで描かれた点線は、おおよそ 2 標準誤差の範囲内です。もし自己相関がこれらの限界内に入っていない場合、おおよそ 5% の有意水準にてゼロとの有意な相違が認められません。正しい ARIMA モデル発見には、練習と経験が必要です。これらの AC, PAC, SC, と AIC は、正しいモデルの仕様を検出するのにとても有用な診断ツールです。

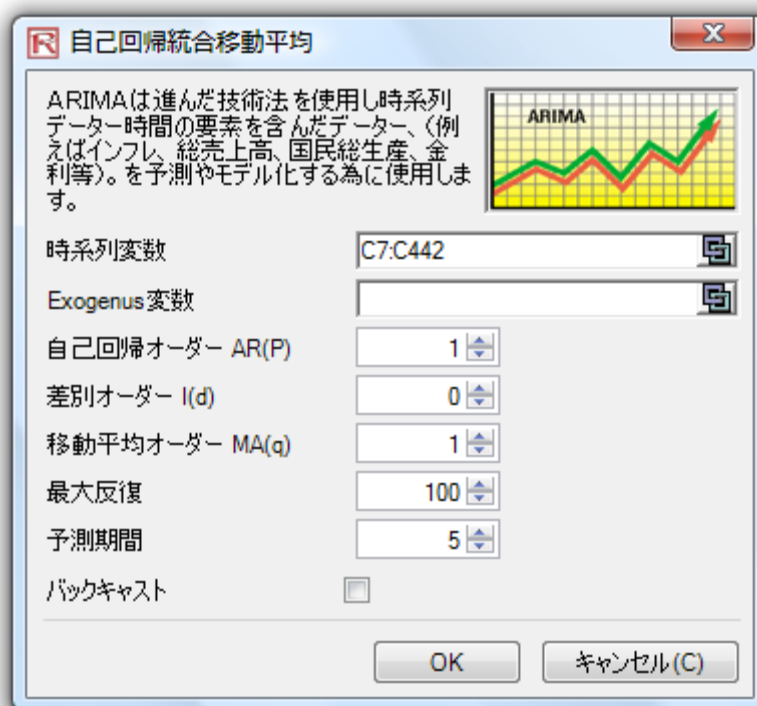


Figure 3.13A Jenkins ボックス ARIMA の予測ツール

ARIMA (Autoregressive Integrated Moving Average)

Regression Statistics

R-Squared (Coefficient of Determination)	0.9999	Akaike Information Criterion (AIC)	4.6213
Adjusted R-Squared	0.9999	Schwarz Criterion (SC)	4.6632
Multiple R (Multiple Correlation Coefficient)	1.0000	Log Likelihood	-1005.1340
Standard Error of the Estimates (SEy)	297.5246	Durbin-Watson (DW) Statistic	1.8588
Number of Observations	435	Number of Iterations	5

Autoregressive Integrated Moving Average or ARIMA(p,d,q) models are the extension of the AR model that use three components for modeling the serial correlation in the time-series data. The first component is the autoregressive (AR) term. The AR(p) model uses the p lags of the time series in the equation. An AR(p) model has the form: $y(t)=a(1)*y(t-1)+...+a(p)*y(t-p)+e(t)$. The second component is the integration (d) order term. Each integration order corresponds to differencing the time series. I(1) means differencing the data once. I(d) means differencing the data d times. The third component is the moving average (MA) term. The MA(q) model uses the q lags of the forecast errors to improve the forecast. An MA(q) model has the form: $y(t)=e(t)+b(1)*e(t-1)+...+b(q)*e(t-q)$. Finally, an ARMA(p,q) model has the combined form: $y(t)=a(1)*y(t-1)+...+a(p)*y(t-p)+e(t)+b(1)*e(t-1)+...+b(q)*e(t-q)$.

The R-Squared, or Coefficient of Determination, indicates the percent variation in the dependent variable that can be explained and accounted for by the independent variables in this regression analysis. However, in a multiple regression, the Adjusted R-Squared takes into account the existence of additional independent variables or regressors and adjusts this R-Squared value to a more accurate view of the regression's explanatory power. However, under some ARIMA modeling circumstances (e.g., with nonconvergence models), the R-Squared tends to be unreliable.

The Multiple Correlation Coefficient (Multiple R) measures the correlation between the actual dependent variable (Y) and the estimated or fitted (Y) based on the regression equation. This correlation is also the square root of the Coefficient of Determination (R-Squared).

The Standard Error of the Estimates (SEy) describes the dispersion of data points above and below the regression line or plane. This value is used as part of the calculation to obtain the confidence interval of the estimates later.

The AIC and SC are often used in model selection. SC imposes a greater penalty for additional coefficients. Generally, the user should select a model with the lowest value of the AIC and SC.

The Durbin-Watson statistic measures the serial correlation in the residuals. Generally, DW less than 2 implies positive serial correlation.

Regression Results

	Intercept	AR(1)	MA(1)
Coefficients	-0.0626	1.0055	0.4936
Standard Error	0.3108	0.0006	0.0420
t-Statistic	-0.2013	1691.1373	11.7633
p-Value	0.8406	0.0000	0.0000
Lower 5%	0.4498	1.0065	0.5628
Upper 95%	-0.5749	1.0046	0.4244

Degrees of Freedom

Degrees of Freedom for Regression	2	Critical t-Statistic (99% confidence with df of 432)	2.5873
Degrees of Freedom for Residual	432	Critical t-Statistic (95% confidence with df of 432)	1.9655
Total Degrees of Freedom	434	Critical t-Statistic (90% confidence with df of 432)	1.6484

Hypothesis Test

The Coefficients provide the estimated regression intercept and slopes. For instance, the coefficients are estimates of the true, population b values in the following regression equation $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + ... + \beta_n X_n$. The Standard Error measures how accurate the predicted Coefficients are, and the t-Statistics are the ratios of each predicted Coefficient to its Standard Error.

The t-Statistic is used in hypothesis testing, where we set the null hypothesis (Ho) such that the real mean of the Coefficient = 0, and the alternate hypothesis (Ha) such that the real mean of the Coefficient is not equal to 0. A t-test is performed and the calculated t-Statistic is compared to the critical values at the relevant Degrees of Freedom for Residual. The t-test is very important as it calculates if each of the coefficients is statistically significant in the presence of the other regressors. This means that the t-test statistically verifies whether a regressor or independent variable should remain in the regression or it should be dropped.

The Coefficient is statistically significant if its calculated t-Statistic exceeds the Critical t-Statistic at the relevant degrees of freedom (df). The three main confidence levels used to test for significance are 90%, 95% and 99%. If a Coefficient's t-Statistic exceeds the Critical level, it is considered statistically significant. Alternatively, the p-Value calculates each t-Statistic's probability of occurrence, which means that the smaller the p-Value, the more significant the Coefficient. The usual significant levels for the p-Value are 0.01, 0.05, and 0.10, corresponding to the 99%, 95%, and 99% confidence levels.

The Coefficients with their p-Values highlighted in blue indicate that they are statistically significant at the 90% confidence or 0.10 alpha level, while those highlighted in red indicate that they are not statistically significant at any other alpha levels.

Analysis of Variance

	Sums of Squares	Mean of Squares	F-Statistic	p-Value	Hypothesis Test	
Regression	38415447.5277	19207723.7638	3171851.1034	0.0000	Critical F-statistic (99% confidence with df of 2 and 432)	4.6546
Residual	2616.0549	6.0557			Critical F-statistic (95% confidence with df of 2 and 432)	3.0166
Total	38418063.5826	19207729.8195			Critical F-statistic (90% confidence with df of 2 and 432)	2.3149

The Analysis of Variance (ANOVA) table provides an F-test of the regression model's overall statistical significance. Instead of looking at individual regressors as in the t-test, the F-test looks at all the estimated Coefficients' statistical properties. The F-Statistic is calculated as the ratio of the Regression's Mean of Squares to the Residual's Mean of Squares. The numerator measures how much of the regression is explained, while the denominator measures how much is unexplained. Hence, the larger the F-Statistic, the more significant the model. The corresponding p-Value is calculated to test the null hypothesis (Ho) where all the Coefficients are simultaneously equal to zero, versus the alternate hypothesis (Ha) that they are all simultaneously different from zero, indicating a significant overall regression model. If the p-Value is smaller than the 0.01, 0.05, or 0.10 alpha significance, then the regression is significant. The same approach can be applied to the F-Statistic by comparing the calculated F-Statistic with the critical F values at various significance levels.

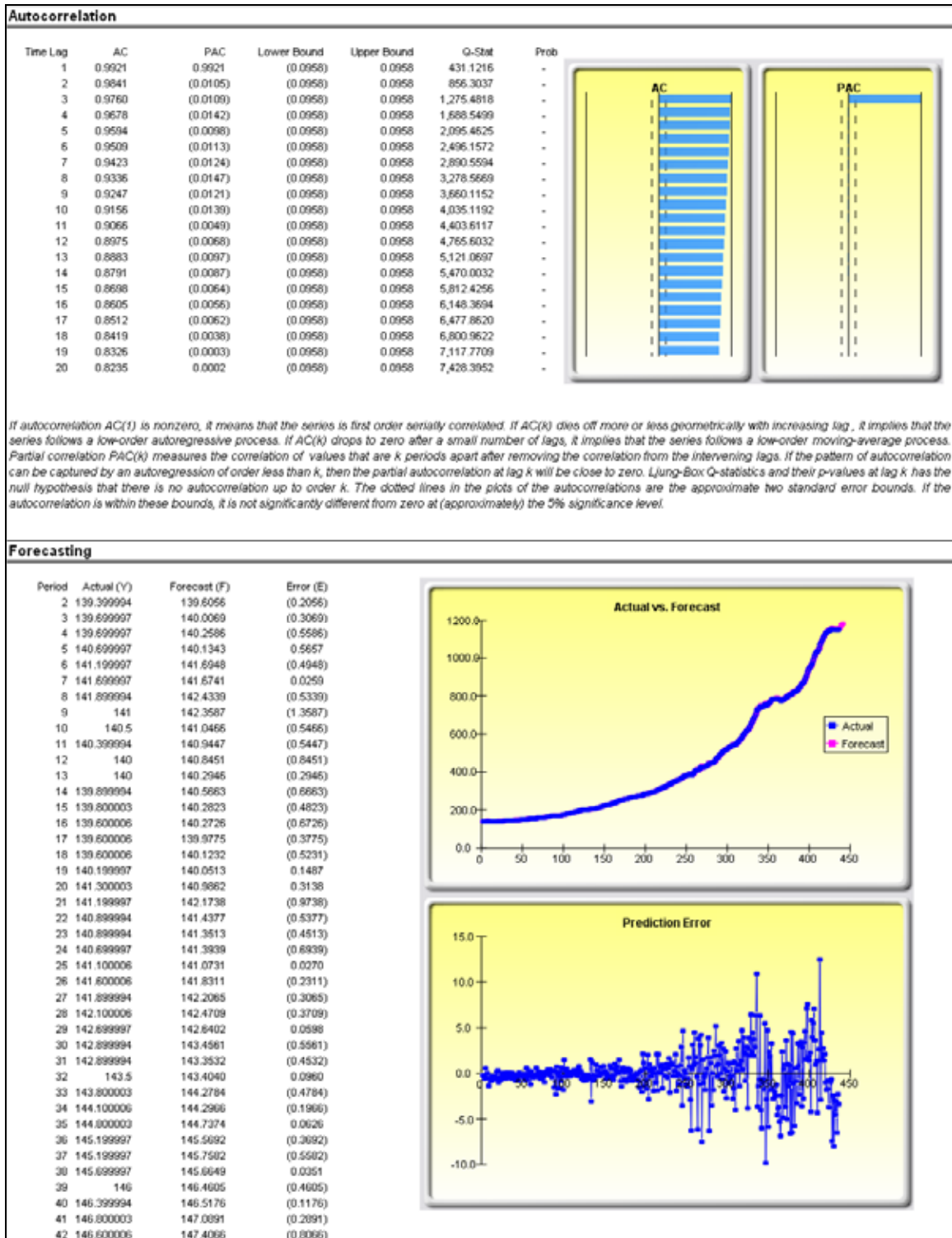


Figure 3.13B Jenkins ボックス ARIMA 予測のレポート

自己 ARIMA (Box-Jenkins ARIMA 高度な時系列)

理論・セオリー:

このツールは、ARIMA モジュールと同じ分析を与えてくれますが、違いとしては、自己・ARIMA モジュールは、モデルの仕様の複数の組み合わせの自動的な検定で幾つかの典型的な ARIMA を自動化し、最適な適合のモデルを戻します。自己・ARIMA の実行は普通の ARIMA 予測に類似しています。違いとしては、P, D, Q の入力は、もはや要求されなく、これらの入力の異なった組み合わせは自動的に実行され、比較されます。

手順:

- Excel を起動し、データを入力するか実在するワークシートと予測する為の履歴データを開いてください (Figure 3.14 で表示されているイラストは、リスクシミュレーターの **例証** メニューにある **高度な予測モデル** のファイルを使用しています)。
- **自己 ARIMA** ワークシートを **リスクシミュレーター (Risk Simulator) | 予測 (Forecasting) | 自己・ARIMA (AUTO-ARIMA)** を選択する事で開いてください。
- リンクアイコンをクリックし、存在する時系列データに関連し、希望する予測周期を入力して **OK** をクリックしてください。



Figure 3.14 AUTO ARIMA モジュール

理論・セオリー:

計量経済学とは、あるビジネス・経済変数の振る舞いのモデル化あるいは予測のためのビジネス分析、モデル化と予測技法の一分野のことを意味しています。基本的な計量経済学のモデルを実行するのは、従属、および独立変数は回帰の実行以前に変換できるという事以外には、普通の回帰分析と同じです。生成されたレポートは重回帰の章で表示されたのと同じで、解釈技法は、以前記述されたものと同じです。

手順:

- Excel をスタートし、データを記入するか、実在するワークシートと予測する為の履歴データを開いてください (Figure 3.15 で表示されたイラストは、リスクシミュレーターの **例証** メニューの **高度な予測モデル** のファイルを使用しています)。
- **基本的な計量経済** のワークシートでデータを選択し、**リスクシミュレーター (Risk Simulator) | 予測(Forecasting) | 基本的な計量経済(Basic Econometrics)** を選択してください。
- 希望する従属、および独立変数を (Figure 3.15 の例証を参照ください) 記入し、**OK** をクリックして、モデルとレポートを実行してください。または、**結果の表示** をクリックして、モデル内で何かの変換を行う場合にはレポートの作成以前に結果を見る事が出来ます。



Figure 3.15 基本的な計量経済学のモジュール

J-S 曲線の予測

理論・セオリー:

J-曲線、または指数成長曲線は、次周期の成長が現在の周期のレベルに依存し、増加が指数関数的である曲線のことです。これは、時間が経つに連れて、一つの周期から次の周期までの値は、かなり増加すると言う事を意味しています。このモデルは、一般的に時間的な変化における生ものの成長と化学反応の予測に使用されます。

手順:

- Excel を起動し、**リスクシミュレーター (Risk Simulator) | 予測(Forecasting) | JS 曲線(JS Curves)**を選択してください。
- J、または S 曲線のタイプを選択し、必要な入力仮定を記入してください (Figures 3.16 と 3.17 を例証として参照してください)。モデルとレポートを実行するのに **OK** をクリックしてください。

J-Curve Exponential Growth Curves

In mathematics, a quantity that grows exponentially is one whose growth rate is always proportional to its current size. Such growth is said to follow an exponential law. This implies that for any exponentially growing quantity, the larger the quantity gets, the faster it grows. But it also implies that the relationship between the size of the dependent variable and its rate of growth is governed by a strict law, of the simplest kind: direct proportion. The general principle behind exponential growth is that the larger a number gets, the faster it grows. Any exponentially growing number will eventually grow larger than any other number which grows at only a constant rate for the same amount of time. This forecast method is also called a J curve due to its shape resembling the letter J. There is no maximum level of this growth curve. Other growth curves include S-curves and Markov Chains.

To generate a J curve forecast, follow the instructions below:

1. Click on Risk Simulator | Forecasting | JS Curves
2. Select Exponential J Curve and enter in the desired inputs (e.g., Starting Value of 100, Growth Rate of 5 percent, End Period of 100)
3. Click OK to run the forecast and spend some time reviewing the forecast report

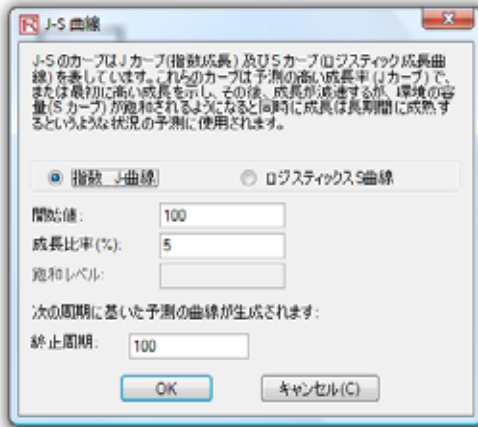


Figure 3.16 J-曲線の予測

S-曲線、またはロジスティック成長曲線は、指数成長率でもって、J-曲線のように開始します。時間が経つに連れ、環境が飽和 (例、市場飽和、競争、混雑など)して、成長は遅くなり、予測値は結局、飽和、または最高レベルで終止してしまいます。このモデルは一般的に、新商品の市場での導入から成熟、低下のサイクルに関するマーケットシェアあるいはの販売成長の予測に使用されます。Figure 3.17 は、サンプル S-曲線を表示しています。

Logistic S Curve

A logistic function or logistic curve models the S-curve of growth of some variable X . The initial stage of growth is approximately exponential, then, as competition arises, the growth slows, and at maturity, growth stops. These functions find applications in a range of fields, from biology to economics. For example, in the development of an embryo, a fertilized ovum splits, and the cell count grows: 1, 2, 4, 8, 16, 32, 64, etc. This is exponential growth. But the fetus can grow only as large as the uterus can hold, thus other factors start slowing down the increase in the cell count, and the rate of growth slows (but the baby is still growing, of course). After a suitable time, the child is born and keeps growing. Ultimately, the cell count is stable; the person's height is constant; the growth has stopped, at maturity. The same principles can be applied to population growth of animals or humans, and the market penetration and revenues of a product, with an initial growth spurt in market penetration, but over time, the growth slows due to competition and eventually the market declines and matures.

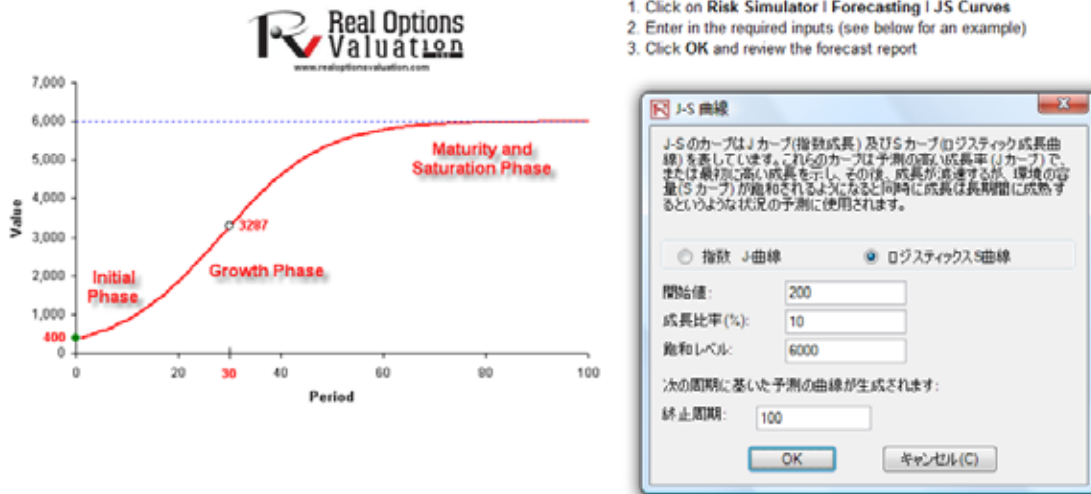


Figure 3.17 S-曲線の予測

GARCH ボラティリティの予測

理論・セオリー:

生成された自己回帰条件付き不均一分散(GARCH) モデルは、そして有価証券の履歴ボラティリティレベルのモデルかや未来のボラティリティレベルの予測（例、株式価格、もの価、石油の価格など）のために使用されます。データの集合は、未加工の価格レベルの時系列でなければいけません。GARCH は、まず価格を相対的なリターンに変換し、履歴データが平均回帰ボラティリティの期間構造に適合するように内部の最適化を実行します。この過程が行われている間、ボラティリティは性質として不均一分散（ある計量経済学の性質に適合した時間的な変化）だと仮定します。GARCH モデルの理論的な詳細は、このユーザーマニュアルの範囲外にあります。GARCH モデルの詳細には、ジョナサン・マン博士の“高度な分析のモデル” (Wiley 2008)をご覧ください。

手順

- Excel を起動し、例証ファイルの **高度な予測のモデル** を選択し、**GARCH** ワークシートを選択し、**リスクシミュレーター(Risk Simulator) | 予測(Forecasting) | GARCH** を選択してください。
- リンクアイコンをクリックし、データの配置を選択し、必要な入力仮定を記入し (Figure 3.18 をご覧ください)、**OK** をクリックして、モデルとレポートを実行してください。

メモ: 一般的なボラティリティの予測に必要な状況は、 $P = 1, Q = 1$, 周期 = 年間の周期数 (月間データには 12、週間データには 52、日次データには、252 または、365)、基礎 = 1 の最小値と周期値の上限と予測周期 = 希望する年次化されたボラティリティの予測の数です。



一般化自己回帰条件付き分散不均一モデル (GARCH)

Historical Data

Days	Inputs
1	459.11
2	460.71
3	460.34
4	460.68
5	460.83
6	461.68
7	461.66
8	461.64
9	465.97
10	469.38
11	470.05
12	469.72
13	466.95
14	464.78
15	465.81
16	465.86
17	467.44
18	468.32
19	470.39
20	468.51
21	470.42
22	470.4
23	472.78
24	478.64
25	481.14
26	480.81
27	481.19
28	480.19
29	481.46
30	481.65
31	482.55
32	484.54
33	485.22
34	481.97
35	482.74
36	485.07

GARCHモデルを実行するには、重要な時系列データを入力し、リスクシミュレーター (Risk Simulator) | 予測 (Forecasting) | GARCH をクリックし、データへの配置のリンクアイコンをクリックし、履歴データの範囲を選択してください (例、C8:C2428)。必要とする項目を入力 (例、P: 1, Q: 1, 年稼働日 252, 予測の基礎 1, 予測期間 10)。OK をクリックしてください。生成された予測のレポートをご覧ください。

GARCH 一般化された自己回帰条件付き分散不均一モデルは、価格自身を使用して金融商品のボラティリティの予測に使用されます。GARCH (P,Q) モデルは、平均回帰式及び実証公式の真なった肯定約 P 及び確約 Q の選択パラメーターを可能にします。肯定約なデータ値だけが GARCH のボラティリティ予測で使用することができることに注意してください。期間とは一年の間にある期間の数を示しています (例えば、12、252、毎日データのための 365)。これらはボラティリティを年間に示すか、または期間約なボラティリティを 1 の期間として保つかです。基礎は、指定された基礎期間です (これは予測の基礎として未来のボラティリティを推定するのにいくつかの過去の期間が必要かと言ふことを表示しています。たとえば 1 から 12 の間の値です。) 実際の目標とは、もしボラティリティの予測を記入した長期間の平均終止時間を與えたい時のことを示します。未加工価格データを年代順 (過去から現在の値を多数の列と単一コラムに表示) に整理することを確かめてください。

データ位置: C8:C2428

生成する GARCH (P,Q) モデル:

P: 1 Q: 1 期間: 252 基礎: 1 予測期間: 10

☒ 目的の可変を適用してください

☒ GARCH ☐ TGARCH ☐ EGARCH ☐ GJR GARCH ☐ GJR TGARCH

OK キャンセル(C)

Figure 3.18 GARCH ボラティリティの予測

	$z_t \sim \text{Normal}$	$z_t \sim T$
GARCH-M	$y_t = c + \lambda \sigma_t^2 + \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$	$y_t = c + \lambda \sigma_t^2 + \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$
GARCH-M	$y_t = c + \lambda \sigma_t + \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$	$y_t = c + \lambda \sigma_t + \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$
GARCH-M	$y_t = c + \lambda \ln(\sigma_t^2) + \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$	$y_t = c + \lambda \ln(\sigma_t^2) + \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$
GARCH	$y_t = x_t \gamma + \varepsilon_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$	$y_t = \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$
EGARCH	$y_t = \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\ln(\sigma_t^2) = \omega + \beta \cdot \ln(\sigma_{t-1}^2) +$ $\alpha \left[\frac{\varepsilon_{t-1}}{\sigma_{t-1}} - E(\varepsilon_t) \right] + r \frac{\varepsilon_{t-1}}{\sigma_{t-1}}$ $E(\varepsilon_t) = \sqrt{\frac{2}{\pi}}$	$y_t = \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\ln(\sigma_t^2) = \omega + \beta \cdot \ln(\sigma_{t-1}^2) +$ $\alpha \left[\frac{\varepsilon_{t-1}}{\sigma_{t-1}} - E(\varepsilon_t) \right] + r \frac{\varepsilon_{t-1}}{\sigma_{t-1}}$ $E(\varepsilon_t) = \frac{2\sqrt{\nu-2} \Gamma((\nu+1)/2)}{(\nu-1)\Gamma(\nu/2)\sqrt{\pi}}$
GJR-GARCH	$y_t = \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 +$ $r \varepsilon_{t-1}^2 d_{t-1} + \beta \sigma_{t-1}^2$ $d_{t-1} = \begin{cases} 1 & \text{if } \varepsilon_{t-1} < 0 \\ 0 & \text{otherwise} \end{cases}$	$y_t = \varepsilon_t$ $\varepsilon_t = \sigma_t z_t$ $\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 +$ $r \varepsilon_{t-1}^2 d_{t-1} + \beta \sigma_{t-1}^2$ $d_{t-1} = \begin{cases} 1 & \text{if } \varepsilon_{t-1} < 0 \\ 0 & \text{otherwise} \end{cases}$

マルコフ連鎖 (Markov Chains)

理論・セオリー:

Markov chain は、未来の状態の確率が、前回の状態に依存し、そのような関係が一緒にリンクされた時に鎖の形状を形成し、長期間的な定常状態のレベルに収斂する時に存在します。このアプローチは一般的に 2 社の競争のマーケットシェアの予測をする時に使用されます。必要とされる入力、最初の店（初期状態）に現れた消費者が次の状態で同じ店に来る確率に対する、競争店に行く遷移確率です。

手順:

- Excel を起動し、**R リスクシミュレーター (Risk Simulator) | 予測(Forecasting) マルコフ チェーン(Markov Chain)**を選択してください。
- 必要な入力仮定を記入し (例証には Figure 3.19 を参照してください)、**OK** をクリックしてモデルとレポートを実行してください。

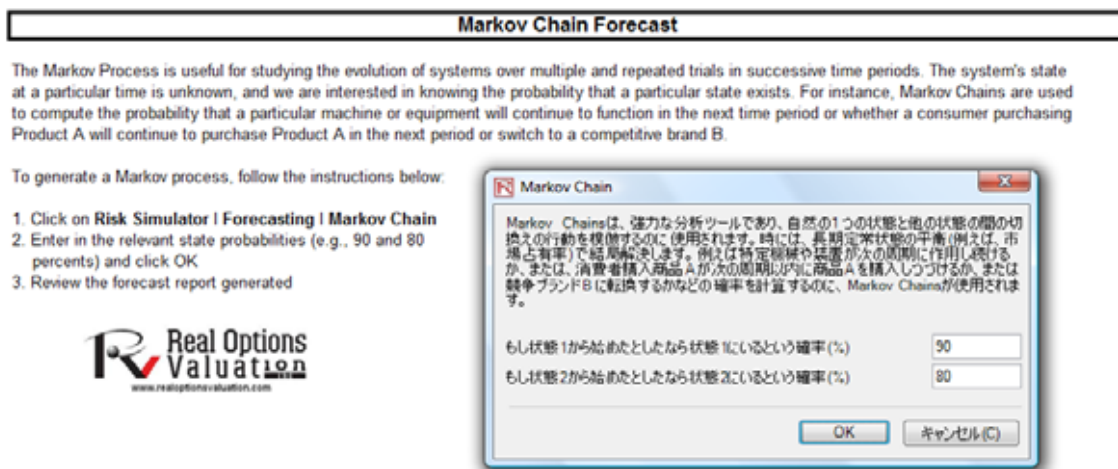


Figure 3.19 Markov Chains 遷移体制)

最尤推定モデル (MLE) : Logit, Probit, Tobit

理論・セオリー:

有限従属変数は、2 値反応（0、1）、切断された、整然された、あるいは、断ち切られたデーターの様なスコープと範囲が有限である従属変数にデーターが含まれている際のシチュエーションを記述します。例えば、従属変数のデーターセットを与えられた

場合(例、年齢、収入、クレジットカード、あるいは抵当貸し付け権者の教育レベル)、最高尤度測定 (MLE)を使用して債務不履行の確立をモデル化することができます。反応、あるいは従属変数 Y が 2 値である場合、1 と 0 の 2 法だけの確立的な結果しかできません(例、 Y は、特定の条件の存在/不在、以前のローンの債務不履行/否・債務不履行、あるデバイスの成功/失敗、概観のはい/いいえの答え等)。また、従属変数回帰 X のベクトルを所持する事で結果 Y を影響すると定義されています。最も一般的な最小 2 乗回帰アプローチは、回帰エラーが不均一分散性であり非・正常である為、無効で、測定された確率測定は、1 以上、0 未満の無意味な値が結果として出ます。MLE 分析は、これらの問題を従属変数が有限である時にログの尤度機能を最大化する為に反復ルーチンの最適化を使用を操作します。

ロジット、あるいはロジスティック回帰は、データーをロジスティック曲線に適合する事でイベントの発生確立を予測する為に使用されます。これは、2 項式回帰にの為に使用される一般化された線形モデルであり、様々な回帰分析の形式と同様に、数的、あるいは分類のどちらかでなければいけない様々な予測変数を利用します。2 値の多重ロジスティック分析で適応された MLE は、特定のグループに属する予測された成功確率を定義する為に従属変数をモデル化するために使用されます。ロジットモデルの為に測定された係数は、ログオッズ確立の比率であり、確立と言う率直な解釈ではありません。初めに簡単な計算を必要とし、アプローチは簡易です。

明確には、ロジットモデルは、測定された $Y = \text{LN}[\text{Pi}/(1-\text{Pi})]$ 、あるいは、逆に言うと、 $\text{Pi} = \text{EXP}(\text{測定された } Y)/(1+\text{EXP}(\text{測定された } Y))$ として定義され、係数 β_i は、ログ確立の比率である為、アンチログ、あるいは、 $\text{EXP}(\beta_i)$ を取る事で、 $\text{Pi}/(1-\text{Pi})$ の確立の比率を得る事ができます。これは、 β_i のユニットを増加するとログの確立の比率もこの値から増加する事を示しています。最後に、確立での変換の比率は、 $dP/dX = \beta_i \text{Pi}(1-\text{Pi})$ です。標準エラーは、予測された係数がどれほど正確なのかを測定し、 t 統計は、これらの標準エラーへの予測された各係数の比率であり、測定された各パラメーターの重要性の一般的な回帰仮説検定の為に使用されます。特定のグループに属する為の成功確率を測定するには(例、一年間に吸った煙草の本数を与えた、喫煙者が肺気管の病気を発生するかどうかの予測)、測定された Y 値を MLE 係数を使用して、計算します。例えば、

$Y=1.1+0.005$ (タバコの本数)がモデルの場合、一年間に100パックを消費する人は、 $1.1+0.005(100)=1.6$ が測定された Y と記述されます。次に、確立の比率のアンチログを計算します。これは、 $\text{EXP}(\text{測定された } Y)/[1+\text{EXP}(\text{測定された } Y)]=\text{EXP}(1.6)/(1+\text{EXP}(1.6))=0.8320$ の様に表示されます。従って、この例証に相当する人物は、83.20%の確立で肺気管の病気を発生する可能性を持っています。

プロビットモデルは(ノーミットモデルとも知られています)、プロビット回帰と呼ばれるアプローチの最高尤度の測定を使用したプロビット機能を利用する2値対応モデルの為の代替りのポピュラーな詳述です。プロビットとロジスティックモデルは、とても相似した予測を生成し、ロジスティック回帰で測定されたパラメーターは、適切なプロビットモデルよりも1.6から1.8回上回りの傾向を表示します。プロビットとロジットのどちらを使用するかは都合により、最も明らかな区別は、ロジスティック分布は、極度な値を計算する為に、高い尖度(ファット・テール)を持っています。例えば、モデル化の対象が家主の判断とし、この反応変数は、2値であり(家の購入、あるいは家を購入しない)、収入、年齢等の従属変数 X_i のシリーズ上に従属し、 $I_i = \beta_0 + \beta_1 X_1 + \dots + \beta_n X_n$ とし、 I_i の値が大きいくほど、家主の確立が高くなります。各家族に、重要な I^* のスレッシュホールドを存在させ、もしも、この値を超えた場合、家が購入され、逆に、超えない場合には、家が購入されません。結果確立(P)は、正規的に分配されていると定義されており、 $P_i = \text{CDF}(I_i)$ の様に標準正規蓄積分布関数(CDF)が使用されています。従って、測定された係数を回帰モデルと測定された Y 値を全く同様に使用し、標準正規分布を適用する事が出来ます(Excel の NORMSDIST 機能、あるいは、リスクシミュレーターの分配的な分析ツールを正規分布を選択し、平均値を0と標準偏差を1と設定し、利用してください)。最後に、プロビット、あるいは、確立的なユニットの測定を得るには、 $I_i + 5$ と設定します(これは、たとえ確立が $P_i < 0.5$ でも、測定される結果が負の数となる為、正規分布は、0の平均値の周辺では対照的になります)。

トビットモデル(断ち切りトビット)は、計量経済学とバイオメトリックスのモデル化技法で、非・負の従属変数 Y_i と1つ、あるいはそれ以上の独立変数 X_i の間の官益を記述する為に使用されます。トビットモデルは、計量経済学モデルであり、従属変数は断ち切られています。従属変数が断ち切られているのは、0より低い値は観測されないから

です。トビットモデルは、隠された観測できない変数 Y^* が存在する事を定義しています。この変数は、 β_i 係数のベクトルを経て X_i 変数上で線形的に従属しており、相互関係を定義します。また、正規的に分配されたエラーの概要 U_i があり、この相互関係上のランダムな影響を取得します。観測できる変数 Y_i は、画されている変数に相当する為に定義されており、たとえ、画されている変数が 0 を上回っているとしても、 Y_i は、0 として定義されます。これは、 $Y_i = Y^*, \text{if } Y^* > 0$ で、 $Y_i = 0, \text{if } Y^* \leq 0$ だからです。 X_i 上で観測できる Y_i の相互関係のパラメーター β_i は、通常の最小 2 乗回帰の使用によって測定され、回帰測定の結果は、矛盾が多く、利回りの下向き方向のバイアスされたスロープ係数と上回りのバイアスされた妨害が表示されます。MLE だけが、トビットモデルに調和し、トビットモデルでは、シグマと呼ばれる補助的な統計が存在し、標準的な最小 2 乗回帰での測定の標準エラーに相当し、測定された係数は、回帰分析と同様に使用されます。

手順:

- Excel を起動し、例証ファイルの **高度な予測モデル**、**MLE** のワークシートを開き、ヘッダーを含めたデータ設定を選択し、**リスクシミュレーター (Risk Simulator) | 予測(Forecasting) | 最高の尤度 (Maximum Likelihood)** を選択してください。
- 下記のリスト(Figure 3.20 を参照)から従属変数を選択し、**OK** をクリックして、モデルとレポートを実行してください。

2項ロジスティック最尤推定法: logit, probit, tobit

LOGIT & PROBIT

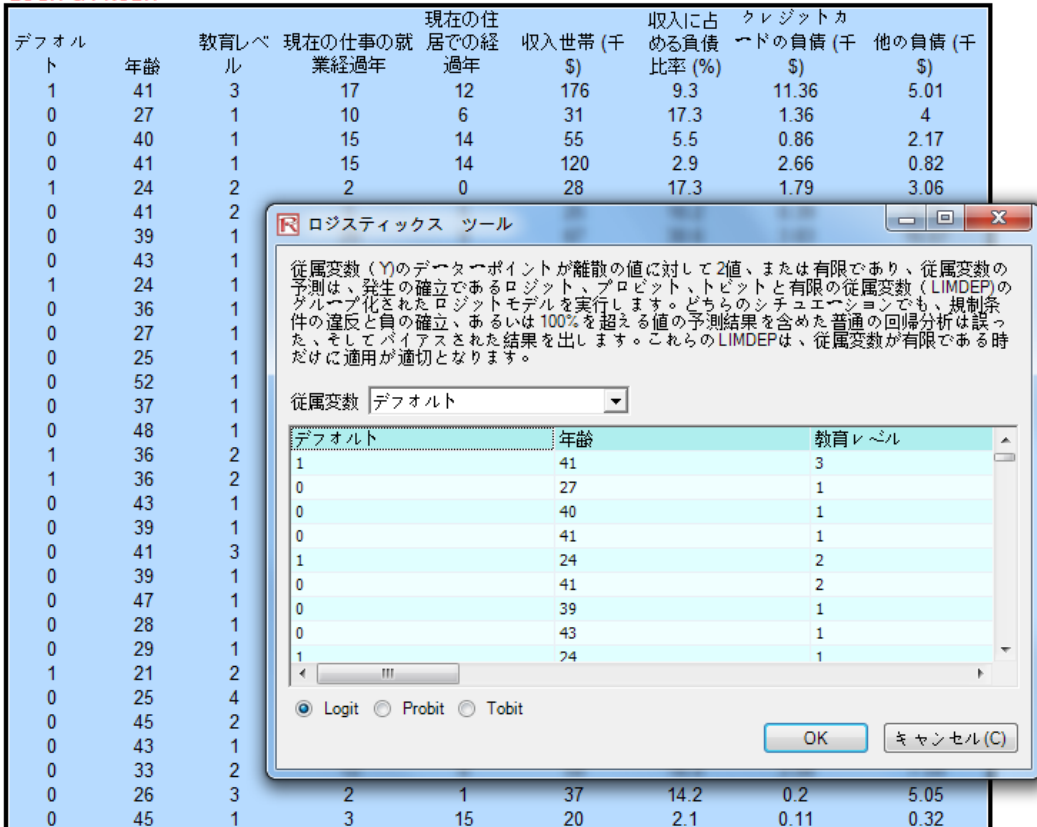


Figure 3.20 最高尤度のモジュール

スプライン (立方スプラインの補入と外押法)

理論・セオリー:

時々、時系列データの設定で欠損値があります。例えば、1年から3年の利率は存在し、5年から8年まで続き、後は10年しかない場合を想定します。スプラインの曲線は、実在するデータに基づいて失われた年の利率を補入する為に使用されます。因みにスプライン曲線は、現在のデータの時間周期を越えて未来の時間周期の値を予測、または外押する為にも使用されます。データは線形、および非線形である事が出来ます。Figure 3.21は、どのようにして立方スプラインが実行されるかを表示しています。知られているX値は、チャートのx-軸上の値（この例証では、知られている年の利率です）を示し、知られているY値は、y-軸上の値（このケースでは知られている利率）を示しています。

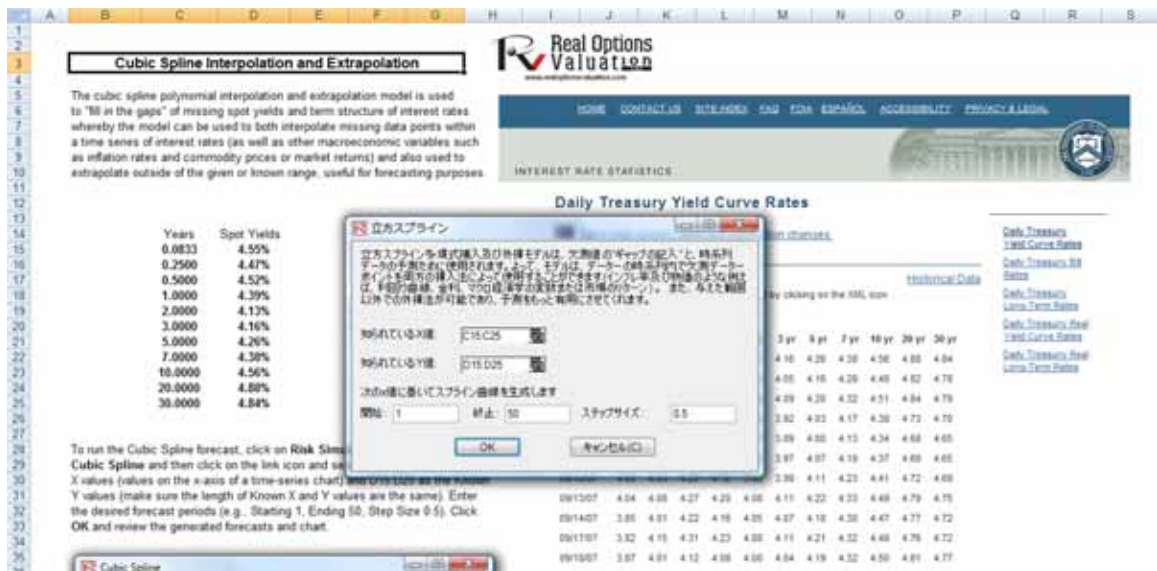


Figure 3.21 立方スプラインモジュール

手順:

- ❖ Excel を起動し、例証ファイルの **高度な予測モデル**、**立方スプライン**ワークシートを開き、ヘッダーを含めたデータ設定を選択し、**リスクシミュレーター (Risk Simulator) | 予測(Forecasting) | 立方スプライン(Cubic Spline)**を選択してください。
- ❖ リンクアイコンをクリックし、知られている X の値と Y の値 (Figure 3.21 を例証として参照してください) をリンクし、補入、および外押の為に必要とする初期値、終止値、これらの値の間のステップワイズを記入し、**OK** をクリックして、モデルとレポートを実行してください。

4. 最適化

この章では、リスクシミュレーターの使用に属する最適化の過程と方法の詳細を参照します。ここでは、静的に対してダイナミック、そして確率的最適化と同様に、連続的にに対して離散的整数の最適化が含まれた方法を使用します。

最適化の方法

幾つかのアルゴリズムは最適化を実行する為に存在し、最適化がモンテカルロ・シミュレーションと共に実行される場合、様々な過程が現れます。リスクシミュレーターでは、三つの異なった最適化の過程と、異なった決定変数のタイプのような最適化のタイプがあります。例えば、リスクシミュレーターは、**整数決定変数** (例、1, 2, 3, 4 または、1.5, 2.5, 3.5 等), **2 値決定変数** (はい・いいえの判断の為に 1 と 0) と、**混合決定変数** (両方とも整数で連続的な変数) と同様に、**連続決定変数** (1.2535, 0.2215 等) を動かす事が出来ます。それ以前に、リスクシミュレーターは、**非線形的な最適化** (例、目的と規制範囲が非線形的な機能と方式の混合と同様に、線形の混合を持っている時) と同様に、**線形的な最適化** (例、目的と制約の両方が全て線形的な方程式か関数の時) をこなす事が出来ます。

最適化プロセスに関する限り、リスクシミュレーターは、**離散的な最適化** を実行するのに使用できます。これは、最適化は、離散的、または静的モデル上で、シミュレーション無しで実行されます。つまり、モデルの全ての入力とは静的で無変換という事です。この最適化のタイプは、不確実性が無く、モデルが知られていると仮定される時に適用されます。因みに、最適なポートフォリオを選び出し、これらに適合する最適な決定変数の割り振りを見出す為に、高度な最適化過程を行う前に、まず離散的な最適化を実行が出来ます。例えば、確率的な最適化問題を実行する以前に、実在する解決があるかを見出す為に、延長された分析を適用する前に離散的な最適化をまず実行します。

次に、最適化がモンテカルロ・シミュレーションと共に適用される時に、**ダイナミックな最適化** を使用します。**シミュレーション・最適化** の名称でも知られています。これは、シミュレーションがまず実行され、この結果は Excel モデルに適用され、最適化は

シミュレーションされた値に適用されます。つまり、シミュレーションは N 試行実行され、その後、最適化の過程は、最適な結果が得られるまで、または実行不可能な設定が出るまで M 回反復されます。これは、リスクシミュレーターの最適化モジュールの使用によってどの予測と仮定統計を使用し、シミュレーションの実行後に置き換えるか、選択が出来ます。その後、予測統計が最適化過程上で適用できます。このアプローチは、複数の相互作用を盛った仮定と予測がある大きなモデルの時に、そして、最適化で幾つかの予測統計が必要とする時の使用が最適です。例えば、仮定、および予測の標準偏差が最適化モデルで必要とされた場合（例、平均をポートフォリオの標準偏差で割った最適化問題と資産の割り振りのシャープレシオ）、このアプローチが使用されます。

一方、確率的最適化の過程は、ダイナミック最適化の過程に似ていますが、違いは全てのダイナミック的な最適化の過程は T 回反復されるという事です。これは、試行回数 N を持ったシミュレーションが実行され、その後、最適な結果を得るために最適化が反復数、 M 実行されます。この後、過程は T 回複製されます。結果として T 値と共に各決定変数の予測チャートが表示されます。つまり、シミュレーションが実行され、予測、または仮定統計は、最適な決定変数の割り振りを見出す為に最適化モデルで使用されます。その後、他のシミュレーションが実行され、異なった予測時計を生成し、これらの更新された値は再び最適化されます。したがって、最終決定変数は、最適な決定変数の範囲を示した各自の予測チャートを持っています。例えば、ダイナミック的な最適化過程で単一・ポイントの推定を得る代わりに、決定変数の分布を得る事が出来、したがって各決定変数の最適な値の範囲を得る事も出来ます。因みに確率過程としても知られています。

最後に、有効フロンティア最適化過程が最適化におけるマージナル増分及びシャドー価格の概念に応用されます。すなわち、制約のどれかがわずかに緩めらる場合、最適化の結果に何が起こるのでしょうか？ 例えて言うと、\$ 1,000,000 として予算の制約が設定されているかどうかです。もしも、制約が \$ 1,500,000、または \$ 2,000,000 の場合、最適な決定とポートフォリオの結果に何が起きるのでしょうか？ これが、投資ファイナンスでの Markowitz の有効的フロンティアで、ポートフォリオの標準偏差が緩やかに増分する場合、どのような付加的なリターンをポートフォリオは生成するのでしょうか？この

過程は、一つの制約を変換することが出来、各変換によってシミュレーションと最適化の過程が実行される事以外には、ダイナミックな最適化の過程に類似しています。この過程は、リスクシミュレーションを使用して手動的に適用するのが最適です。これは、ダイナミック、および確率的最適化を実行する事で、その後、制約と共に他の最適化を再実行し、この過程を何回か反復するという事です。この手動的な過程は、付加的な分析の価値があるか、または目的と決定変数の有意な変換を得る為に、制約での最低限の増加がどれだけ離れていなければいけないか、制約の変換によって結果が類似しているか、相違しているかを定める分析の様に重要です。

一つの項目は考察の価値があります。確率的な最適化を実行する他のソフトウェアがありますが、実用的ではありません。例えば、シミュレーションの実行後、最適化過程の一回の反復が生成され、他のシミュレーションが実行された後、2回目の最適化の反復が実行される為、時間と資源の無駄です。これは、最適化では、最適な結果を得る為に、複数の反復（複数から千の反復を及ぶ）が必要とされる厳格なアルゴリズムの設定を通してモデルに導入されます。したがって、一つの反復の生成は、時間と資源の無駄です。同じポートフォリオは、何時間も必要とする前に記述されたアプローチを使用するよりも、リスクシミュレーターを通して数分で解決する事が出来ます。シミュレーション・最適化のアプローチは一般的に悪い結果を齎すけれども、確率的最適化のアプローチではそうではありません。モデルに最適化を適用する際には、最大の注意を忘れないでください。

次に二つの最適化の問題があります。一つは、連続的な決定変数を使用し、もう一つは、離散的な整数の決定変数を使用します。どちらかのモデルで、離散的最適化、ダイナミックな最適化、確率的な最適化が適用出来、および有効フロンティアと影の価格が適用できます。これらのアプローチのどれかが二つの例証に使用できます。したがって、簡潔には、モデルの設定だけが表示され、ユーザーの為にどの最適化を使用するかを選択が可能です。また、連続的なモデルは、整数の最適化の二つ目の例証が線形的な最適化モデルの例証（これらの目的と全ての制約は線形的）であるのに対して、非線形的な最適化のアプローチ（これは、計算されたポートフォリオのリスクは非線形的な公式を持ち、目的は、ポートフォリオのリターンを非線形的な公式でポートフォリオのリスク

で割ったものです)を使用します。したがって、これらの二つの例証は、前に記述した全ての過程が含まれています。

最適化と連続的な決定変数

Figure 4.1 は、連続的な最適化モデルのサンプルを表示しています。ここで使用されている例証は、**連続的な最適化**を使用し、このファイルはスタートメニューにあります。

スタート (Start) | リアルオプションズバリュエーション (Real Options Valuation) | リスクシミュレーター (Risk Simulator) | 例証 (Examples)、または直接 **リスクシミュレーター (Risk Simulator) | 例証モデル (Example Models)**を通してアクセスしてください。この例証では、10 の異なった資産クラス（例、様々なタイプの投資信託、株式、および資産）があり、アイデアとしては、最も有効な割り振りを持ったポートフォリオは最も最適な利益を得るという事です。これは、可能な最適なポートフォリオのリターンを生成する為には、各資産クラスに固有リスクを与えます。最適化のコンセプトを解釈する為に、どのようにして最適化の過程がより良い方法で適用できるかを見る為にサンプルモデルをもっと深く勉強しなければいけません。モデルは、10 の資産クラスを表示し、これらの各資産は独自の年間リターンと年間ボラティリティを持っています。これらのリターンとリスクの測定は、異なった資産クラス同士で一貫的に比較できるように年間毎に表示されています。リターンは、比較可能なリターンの幾何学的平均を使用して計算されているのに対して、リスクは、対数的に比較可能な株式リターンのアプローチを使用して計算されています。株式、および資産クラス上での年間リターンと年間ボラティリティの計算の詳細には、この章の付録を参照してください。

り、 $R_p = \omega_A R_A + \omega_B R_B + \omega_C R_C + \omega_D R_D$ で、 R_p は、ポートフォリオ上のリターンを示し、 $R_{A,B,C,D}$ は、プロジェクトの個別のリターンを示し、 $\omega_{A,B,C,D}$ は、各ウエイトか各プロジェクトの資本の割り振りです。

また、セル D17 で多角化されたリスクのポートフォリオは、

$$\sigma_p = \sqrt{\sum_{i=1}^i \omega_i^2 \sigma_i^2 + \sum_{i=1}^n \sum_{j=1}^m 2\omega_i \omega_j \rho_{i,j} \sigma_i \sigma_j}$$

によって計算されています。ここでは ρ_{ij} は、

資産クラスの中の相互相関関係に当たり、したがって、相互相関が負の場合、リスクの多様化の効果が見られ、ポートフォリオのリスクは減少します。但し、計算を簡単にする為に、このポートフォリオのリスクの計算を通して、資産クラスの中の相関をゼロと考慮しますが、先ほど記述したように、リターンのシミュレーションの適用の際には、相関を考慮してください。したがって、これらの異なった試算リターンの中で静的相関を適用する代わりに、これら自身のシミュレーション仮定に相関を適用し、シミュレーションされたリターンの値の間でもっとダイナミックな関係が現れます。

最後に、リスクのリターンへの比率、およびシャープレシオは、ポートフォリオのために計算されます。この値は、セルの C18 で表示され、この最適化の実行で最大化させる目的を表しています。まとめとして、この例証モデルの詳細が下記に記述されています。

目的:	リスクへの最大のリターン比率(C18)
決定変数:	割り振りのウエイト (E6:E15)
決定変数の制限:	必要とされる最小値と最大値 (F6:G15)
制約:	総合的な割り振りの足し算は 100% (E17)

手順:

- 📁 例証ファイルを開き、新規プロファイルを **リスクシミュレーター (Risk Simulator) | 新規プロファイル (New Profile)** クリックして起動し、名称をつけてください。
- 📁 最適化の最初のステップとして決定変数の設定を行ってください。セルの E6 を選択し、最初の決定変数を設定 (**リスクシミュレーター (Risk Simulator) | 最適**

化 (Optimization) | 決定の設定 (Set Decision))をし、リンクアイコンをクリックして、セル F6 と G6 の下限と上限の値のように名称セル(B6)を選択してください。その後、リスクシミュレーターのコピーを使用して、セル E6 の決定変数をコピーし、残ったセルの E7 から E15 に貼り付けてください。

■ 最適化の二つ目のステップとして、制約を設定する事です。ここでは一つの制約しかありません。これは、ポートフォリオの総合的な割り振りの足し算は 100% にならなければならないという事を示しています。その為には、**リスクシミュレーター (Risk Simulator) | 最適化 (Optimization) | 制約 (Constraints) ...** そして、**追加**を選択し、新しい制約を加えてください。その後、セル E17 を選択し、100%に同一(=)させてください。終了後、OK をクリックしてください。

■ 最適化の最後のステップとして、目的関数の設定と目的のセル C18 と **リスクシミュレーター (Risk Simulator) | 最適化 (Optimization) | 最適化の実行 (Run Optimization)** の選択を通して最適化を実行させ、最適化（静的最適化、ダイナミック的最適化、および確率的最適化）の選択をして下さい。起動するにあたって、**静的最適化**を選択してください。目的セルは C18 へ設定されているかどうかを確認した後、最大化を選択してください。必要であれば、制約と決定変数を検討する事が出来ます。または、静的最適化を実行するにあたって OK をクリックしてください。

■ 最適化が一度完了した後、目的と同様に、決定変数を元の値に戻す為に、**戻す**を選択するか、最適化された決定変数を適用する為に、**置き換える**を選択してください。一般的に、置き換えるは、最適化の実行後に選択されます。

Figure 4.2 は、これらの上記の過程のステップを表示しています。モデル上のリターンとリスク(コラム C と D)にシミュレーション仮定を追加する事が出来、ダイナミックな最適化と確率的最適化を付加の練習として適用する事が出来ます。

決定変数設定

決定名称(N)

決定タイプ

☒ 連続(例え 1.15, 2.35, 10.55)(C)

下回る 上回る

☐ 整数(例え 1, 2, 3)(I)

下回る 上回る

☐ 2値(0または1)(B)

OK キャンセル(C)

規制

現在の規制

☒ \$E\$17 == 100%

追加(A)
変換(H)
削除(D)
OK
キャンセル(C)

最適化の目的

目的セル(O)

最適化の目的

☒ 目的セルの値を最大化(A)

☐ 目的セルの値を最小化(I)

OK キャンセル(C)

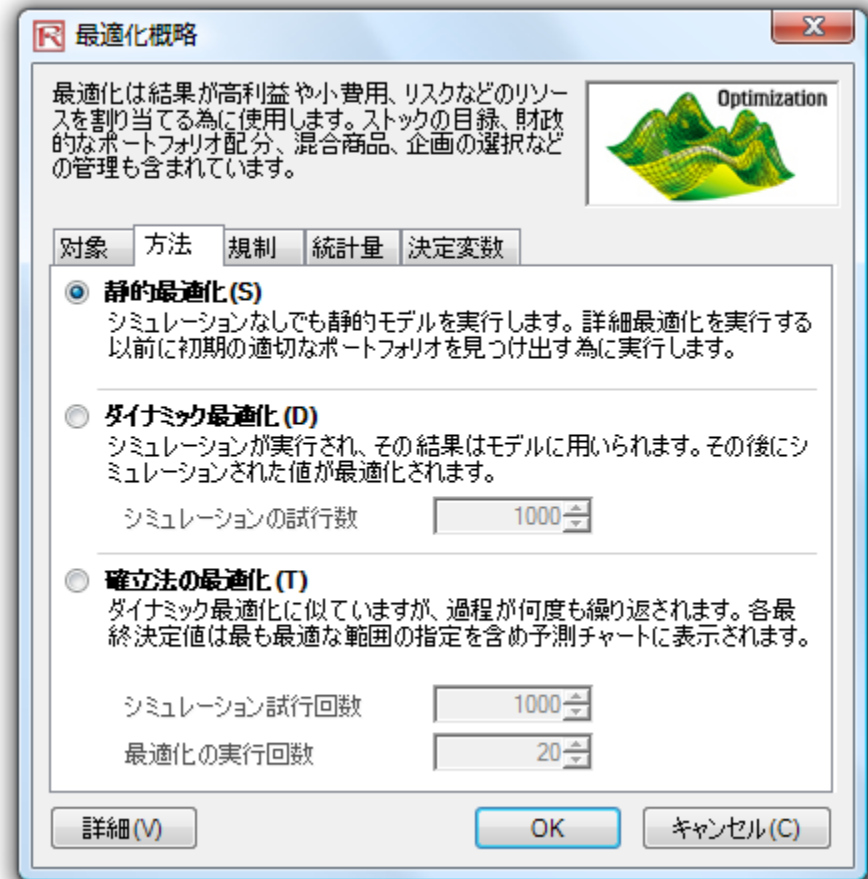


Figure 4.2 リスク シミュレーターで連続最適化の実行

結果の解釈:

最適化の最終結果は Figure 4.3 に表示されており、ポートフォリオの為の資産の最適な割り振りはセル E6:E15 にあります。これは、各資産の変動が 5%と 35%の間で、割り振りの足し算が 100%でなければいけないと言う制限を与えられ、リスクへのリターン比率の最大化は Figure 4.3 に表示されています。

これまでの結果と最適化の仮定の適用の検討で多少重要な事はメモしなければいけません。

- 正しい最適化の実行は利益、およびリスクに対するリターンシャープレシオのを最大化することです。
- もしも、その代りに総合ポートフォリオのリターンを最大化した場合、最適な割り振りの結果は乏しく、取得に最適化は必要ありません。これは、低い 8 つの資産に対して割り振り 5% (認められた最小値)、最もリターンが高い資産に対して 35%(認められた最大値)、残り(25%)は、2 番目に良いリターンを持った資産に割り当てます。最適化は必要ではありません。但し、このような方法でポートフォリオを割り振りする時は、ポートフォリオのリターン自身も高くなりますが、リスクへのリターン比率を最大化する時に比べてリスクが大いに高くなってしまいます。
- 一方、総合的なポートフォリオリスクを減少する事が出来ませんが、リターンも低くなります。

Table 4.1 は、最適化された三つの異なった目的から成った結果を表示しています。

目的:	ポートフ ォリオの リターン	ポートフ ォリオの リスク	ポートフォ リオのリス クへのリタ ーン比率
リスクへのリターン比率の最大 化	12.69%	4.52%	2.8091
リターンの最大化	13.97%	6.77%	2.0636
リスクの最小化	12.38%	4.46%	2.7754

Table 4.1 最適化の結果

表から、最適なアプローチはリスクへのリターン比率を最大化することであり、これは、同じ量のリスクには、この割り振りは最も高いリターンを齎してくれます。一方、同じ量のリターンには、この割り振りは最も低いリスクの確率を与えてくれます。この利益、およびリスクへのリターン比率のアプローチは、モダンなポートフォリオ理論での Markowitz の有効フロンティアの礎石です。これは、総合的なポートフォリオのリスクレベルを制約し、時間が経つに連れてこれを増加する事で、異なったリスクの性質の為の様々な有効なポートフォリオの割り振りを得る事が出来ます。したがって、異なった有効ポートフォリオの配分は、異なったリスク選好を持つ異なった個々人の為に得る事が出来ます。

資産配分最適化モデル

資産クラスの記述	年間収益率	ボラティリティ リスク	配分の重量	最小必要な 配分	最高必要な 配分	リスクの比率 への収益	リターン ランキン グ (Hi-Lo)	リスク キング (Lo- Hi)	ラン グへのリター ン (Hi-Lo)	配分ランキン グ (Hi-Lo)
資産クラス 1	10.54%	12.36%	11.09%	5.00%	35.00%	0.8524	9	2	7	4
資産クラス 2	11.25%	16.23%	6.86%	5.00%	35.00%	0.6929	7	8	10	10
資産クラス 3	11.84%	15.64%	7.78%	5.00%	35.00%	0.7570	6	7	9	9
資産クラス 4	10.64%	12.35%	11.23%	5.00%	35.00%	0.8615	8	1	5	3
資産クラス 5	13.25%	13.28%	12.09%	5.00%	35.00%	0.9977	5	4	2	2
資産クラス 6	14.21%	14.39%	11.04%	5.00%	35.00%	0.9875	3	6	3	5
資産クラス 7	15.53%	14.25%	12.30%	5.00%	35.00%	1.0898	1	5	1	1
資産クラス 8	14.95%	16.44%	8.90%	5.00%	35.00%	0.9094	2	9	4	7
資産クラス 9	14.16%	16.50%	8.37%	5.00%	35.00%	0.8584	4	10	6	8
資産クラス 10	10.06%	12.50%	10.35%	5.00%	35.00%	0.8045	10	3	8	6
ポートフォリオ 合計 リスク率への収益	12.6919%	4.52%	100.00%							

Figure 4.3 連続最適化の結果

最適化と離散的整数の変数

時々、決定変数は連続的ではありませんが、離散的整数 (例、0 と 1) です。これは、このような最適化を on-off のスイッチ、または go/no-go の判断で適用できるという事です。Figure 4.4 は、20 のプロジェクトが表示された中からのプロジェクトの選択モデルを表示しています。ここでの例証は、**離散的な最適化**ファイルを使用し、スタートメニューから **スタート (Start) | リアルオプションズバリュエーション (Real Options Valuation) | リスクシミュレーター (Risk Simulator) | 例証 (Examples)** を選択するか、**リスクシミュレーター (Risk Simulator) | 例証モデル (Example Models)** をクリックして直接開いてください。前の様な各プロジェクトは、独自のリターン (拡張された正味現在価値と正味現在価値

の為の ENPV と NPV を持っており、ENPV は、単に NPV にある戦略的なリアルオプションの値が足されたもの(の)で、実施の費用、リスク等を有しています。必要な場合、このモデルは必要とされたフル・タイムの等価値(FTE)、他の様々な機能の資源、そしてこれらの付加の資源に設定できる制約を含むように修正できます。このモデルの中の入力は一般的に、他のスプレッドシートモデルによってリンクされています。例えば、各プロジェクトは、独自の割引のキャッシュフロー、および資産モデルのリターンを持っています。ここでのアプリケーションは、ある予算配分の制約においてポートフォリオのシャープレシオを最大化することです。このモデルの様々なバージョンが作成できます。例えば、ポートフォリオのリターンを最大化するか、リスクを最小するか、選択するプロジェクトの総合的な数は6を超えてはいけないなどの制約を追加します。これらの全てのアイテムは、存在するモデルを使用して実行する事が出来ます。

手順:

- 例証ファイルを開き、**リスクシミュレーター (Risk Simulator) | 新規プロファイル (New Profile)** をクリックする事で新規プロファイルを起動し、名称を与えてください。
- 最適化の最初のステップとして決定変数の設定を行ってください。セルの J4 を選択し、最初の決定変数を設定 (**リスクシミュレーター (Risk Simulator) | 最適化 (Optimization) | 決定の設定 (Set Decision)**)をし、リンクアイコンをクリックして、名称セル(B4)を付け、**2 値**の変数を選択してください。その後、リスクシミュレーターのコピーを使用して、セル J4 の決定変数をコピーし、残ったセルの J5 から J15 に貼り付けてください。これが、幾つかの決定変数しか持っていない時と、各決定変数にユニークな、独自の名称を付ける事が出来る場合に最適な方法です。
- 最適化の二つ目のステップとして、制約を設定する事です。ここでは二つの制約があります。これは、ポートフォリオの総合的な割り振りの予算は、\$ 5,000 以上でなければいけなく、プロジェクトの総合数は6を超えてはいけないと言う事を示しています。その為には、**リスクシミュレーター (Risk Simulator) | 最適化 (Optimization) | 制約 (Constraints) ...** そして、**追加**を選択し、新しい制約

を加えてください。その後、セル D17 を選択し、5,000 に同一、またはより小さく($\leq$)させてください。セル J17 ≤ 6 と設定した後、繰り返してください。

- 最適化の最後のステップとして、目的関数の設定と目的のセル C19（または C17）と **リスクシミュレーター (Risk Simulator) | 最適化 (Optimization) | 目的の設定 (Set Objective)** の選択を通して最適化を実行させ、**リスクシミュレーター (Risk Simulator) | 最適化 (Optimization) | 最適化の実行 (Run Optimization)** を選択し、最適化（静的最適化、ダイナミック的最適化、および確率的最適化）の選択をして下さい。起動するにあたって、**静的最適化**を選択してください。目的セルは、シャープレシオかポートフォリオのリスクへのリターン比率かどうかを確認した後、最大化を選択してください。必要であれば、制約と決定変数を検討する事が出来ます。または、静的最適化を実行するにあたって OK をクリックしてください。

Figure 4.5 は、前に記述された過程を表示しています。モデルの ENPV とリスク（コラム C と F)上にシミュレーション仮定を追加する事が出来、付加の練習の為には、ダイナミックな最適化と確率的な最適化を適用できます。

	A	B	C	D	E	F	G	H	I	J
1										
2										
3		クレジットライン	ENPV	費用	リスク	リスク	リスク率へのリターン	確立の索引		セクション
4		プロジェクト 1	\$458.00	\$1,732.44	\$54.96	12.00%	8.33	1.26		1.00
5		プロジェクト 2	\$1,954.00	\$859.00	\$1,914.92	98.00%	1.02	3.27		1.00
6		プロジェクト 3	\$1,599.00	\$1,845.00	\$1,551.03	97.00%	1.03	1.87		1.00
7		プロジェクト 4	\$2,251.00	\$1,645.00	\$1,012.95	45.00%	2.22	2.37		1.00
8		プロジェクト 5	\$849.00	\$458.00	\$925.41	109.00%	0.92	2.85		1.00
9		プロジェクト 6	\$758.00	\$52.00	\$560.92	74.00%	1.35	15.58		1.00
10		プロジェクト 7	\$2,845.00	\$758.00	\$5,633.10	198.00%	0.51	4.75		1.00
11		プロジェクト 8	\$1,235.00	\$115.00	\$926.25	75.00%	1.33	11.74		1.00
12		プロジェクト 9	\$1,945.00	\$125.00	\$2,100.60	108.00%	0.93	16.56		1.00
13		プロジェクト 10	\$2,250.00	\$458.00	\$1,912.50	85.00%	1.18	5.91		1.00
14		プロジェクト 11	\$549.00	\$45.00	\$263.52	48.00%	2.08	13.20		1.00
15		プロジェクト 12	\$525.00	\$105.00	\$309.75	59.00%	1.69	6.00		1.00
16										
17		合計	\$17,218.00	\$8,197.44	\$7,007	40.70%				12.00
18		目的:	最大	<=\$5000						<=6
19		シャープの比率	2.46							
20										
21		ENPVは各クレジット・ラインのおよびプロジェクトの予期されたNPVであり、その間、費用はクレジット・ラインをカバーする為に								
22		必須の重要な保有物と同様で、管理の総額とも考えられます。リスクはクレジット・ラインのENPVの変化の係数です。								
23										

Figure 4.4 離散系整数の最適化モデル

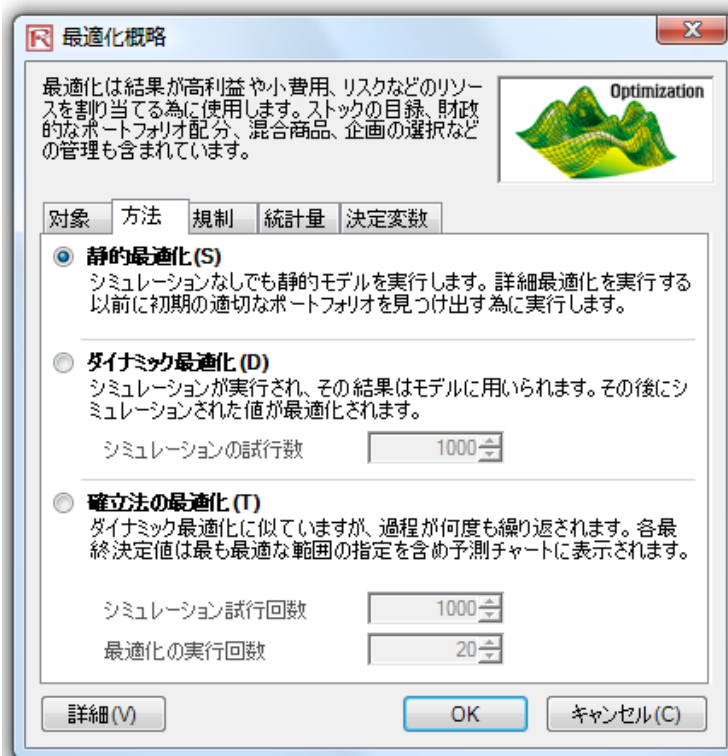
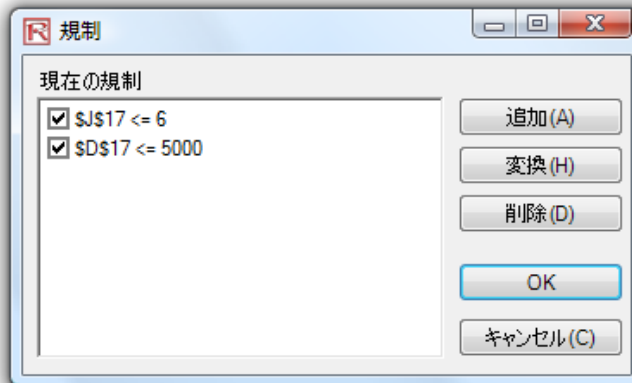


Figure 4.5 リスクシミュレーターで離散系整数の最適化の実行

結果の解釈:

Figure 4.6 は、シャープレシオを最大化するプロジェクトの最適な選択のサンプルを表示しています。一方、一つは常に総合収入を最大化できますが、これは前述のように乏しい過程で、最も高いリターンのプロジェクトを選ぶ事を単に含み、制約予算を超えるまで、またはお金を使い果たすまでリストを辿って行きます。この実行は、高い利益を齎すプロジェクトは一般的に高いリスクを持っている様に、論理的に望ましくないプロジェクトを齎します。希望の場合、リスクの値と ENPV に仮定を追加する事で、確率、またはダイナミック的な最適化を使用する事で最適化を複製する事が出来ます。

クレジットライン	ENPV	費用	リスク	リスク	リスク率へのリターン	確立の索引	セレクション
プロジェクト 1	\$458.00	\$1,732.44	\$54.96	12.00%	8.33	1.26	1.00
プロジェクト 2	\$1,954.00	\$859.00	\$1,914.92	98.00%	1.02	3.27	0.00
プロジェクト 3	\$1,599.00	\$1,845.00	\$1,551.03	97.00%	1.03	1.87	0.00
プロジェクト 4	\$2,251.00	\$1,645.00	\$1,012.95	45.00%	2.22	2.37	1.00
プロジェクト 5	\$849.00	\$458.00	\$925.41	109.00%	0.92	2.85	0.00
プロジェクト 6	\$758.00	\$52.00	\$560.92	74.00%	1.35	15.58	1.00
プロジェクト 7	\$2,845.00	\$758.00	\$5,633.10	198.00%	0.51	4.75	0.00
プロジェクト 8	\$1,235.00	\$115.00	\$926.25	75.00%	1.33	11.74	1.00
プロジェクト 9	\$1,945.00	\$125.00	\$2,100.60	108.00%	0.93	16.56	0.00
プロジェクト 10	\$2,250.00	\$458.00	\$1,912.50	85.00%	1.18	5.91	0.00
プロジェクト 11	\$549.00	\$45.00	\$263.52	48.00%	2.08	13.20	1.00
プロジェクト 12	\$525.00	\$105.00	\$309.75	59.00%	1.69	6.00	1.00
合計	\$5,776.00	\$3,694.44	\$1,539	26.64%			6.00
目的:	最大	<=\$5000					<=6
シャープの比率	3.75						

ENPV は各クレジット・ラインのおよびプロジェクトの予期されたNPVであり、その間、費用はクレジット・ラインをカバーする為に必須の重要な保有物と同様で、管理の総額とも考えられます。リスクはクレジット・ラインのENPVの変化の係数です。

Figure 4.6 シャープ比を最大にする最適なプロジェクトの選択

活動している最適化の付加の例証には、積分されたリスク分析の第 11 章のケーススタディのリアルオプションズ分析: ツールと技法、第2版 (Wiley Finance, 2005)を参照して下さい。このケースは、どのようにして有効なフロンティアが生成できるか、どのようにして予測、シミュレーション、最適化とリアルオプションが継ぎ目の無い分析過程の中で混合できるかを表示しています。

5. リスクシミュレーション分析ツール

この章は、リスクシミュレーターの分析ツールと一緒に取引をします。この分析ツールはリスクシミュレーターのソフトウェアの例証のアプリケーションを通して議論します。ステップ事に表示されたイラストと共に完了してください。リスク分析分野での分析家の仕事にとってこのツールはとっても貴重とされています。適用可能な各ツールはこの章で細かく記述されています。

シミュレーションでの竜巻と感度ツール

理論・セオリー:

一つの強力なシミュレーションツールは竜巻分析で、モデルの結果への各変数の静的な影響を獲得します。これは、ツールはモデル内の各変数を初期値に自動的に攪乱させ、モデルの予測上の変動を取得し、結果として成った攪乱を一番有意なものからそうでないものの順にランクします。Figures 5.1 から 5.6 は、竜巻分析のアプリケーションを表示しています。例えば、Figure 5.1 は、シンプルな割り引いたキャッシュフローのモデルで、モデルの入力仮定は表示されています。ここでの疑問は、モデルの結果に最も影響するクリティカルな成功ドライバーはどれか?と言う事です。これは、どれだけ本当に\$96.63 の正味現在価値を動かしているのか、または、どの入力変数が一番値に影響を与えるか?と言う事です。

竜巻チャートのツールは、**リスクシミュレーター (Risk Simulator) | ツール (Tools) | 竜巻分析 (Tornado Analysis)** を通して得る事が出来ます。また、最初の例証に続いて、例証フォルダーから **竜巻と感度チャート (線形)** ファイルを開きます。Figure 5.2 は、正味現在価値が含まれたセル G6 が分析される目的の結果として選択されるサンプルモデルが表示されています。モデルの目的セルの先例は、竜巻チャートを作成する為に使用されます。先例とは、モデルの結果に影響を及ぼす全ての入力と中間の変数です。例えば、モデルが $A = B + C$ で、および $C = D + E$ で、B, D, と E は A の先例(C は、先例ではなく中間の計算された値の場合)です。Figure 5.2 は、目的の結果の推定の為に、各先例の変数のテスト範囲を表示しています。もしも、先例の変数がシンプルな入力の場合、選択された範囲（例、デフォルトは±10%）に基づいた検定の範囲は、シンプルな混乱です。各先例の変数は、必要な場合、様々な百分位数で攪乱できます。予期された値の周囲の小さい混乱よりもむしろ極値の方が検定しやすい様に幅広い範囲は重要です。特定

	A	B	C	D	E	F	G	H	I	J
1										
2										
3										
4		基年の年	2005		合計現在価値の純利点	\$1,896.63				
5		マーケットリスク割引が調節された比率	15.00%		合計現在価値の投資	\$1,800.00				
6		プライベート・リスクの割引率	5.00%		純現在価値	\$96.63				
7		年間販売の成長率	2.00%		内部収益率	18.80%				
8		価格の腐食率	5.00%		投資の収益	5.37%				
9		有効な税率	40.00%							
10										
11										
12										
13										
14										
15										
16										
17										
18										
19										
20										
21										
22										
23										
24										
25										
26										
27										
28										
29										
30										
31										
32										
33										
34										
35										
36										
37										
38										
39										
40										
41										
42										

割引かれた現金流動モデル

既定の設定

予測法の設定

	2005	2006	2007	2008	2009
商品 A 標準価格	\$10.00	\$9.50	\$9.03	\$8.57	\$8.15
商品 B 標準価格	\$12.25	\$11.64	\$11.06	\$10.50	\$9.98
商品 C 標準価格	\$15.15	\$14.39	\$13.67	\$12.99	\$12.34
商品 A 量	50.00	51.00	52.02	53.06	54.12
商品 B 量	35.00	35.70	36.41	37.14	37.89
商品 C 量	20.00	20.40	20.81	21.22	21.65
合計収入	\$1,231.75	\$1,193.57	\$1,156.57	\$1,120.71	\$1,085.97
売り上げ費用	\$104.76	\$179.03	\$173.40	\$168.11	\$162.90
売り上げ総利益	\$1,046.99	\$1,014.53	\$983.08	\$952.60	\$923.07
検査費	\$157.50	\$160.65	\$163.86	\$167.14	\$170.48
SG&A 費用	\$15.75	\$16.07	\$16.39	\$16.71	\$17.05
検査収入 (EBITDA)	\$873.74	\$837.82	\$802.83	\$768.75	\$735.54
下落	\$10.00	\$10.00	\$10.00	\$10.00	\$10.00
償却	\$3.00	\$3.00	\$3.00	\$3.00	\$3.00
EBIT	\$860.74	\$824.82	\$789.83	\$755.75	\$722.54
利払い	\$2.00	\$2.00	\$2.00	\$2.00	\$2.00
EBT	\$858.74	\$822.82	\$787.83	\$753.75	\$720.54
税金	\$343.50	\$329.13	\$315.13	\$301.50	\$288.22
純収益	\$515.24	\$493.69	\$472.70	\$452.25	\$432.32
下落	\$13.00	\$13.00	\$13.00	\$13.00	\$13.00
純流動比率の変更	\$0.00	\$0.00	\$0.00	\$0.00	\$0.00
資本支出	\$0.00				

© 2005-2012 Real Options Valuation, Inc.

手順

- Excel モデル(例、この例証ではセルの G6 が選択されています)で単一の結果セル(例、公式、または方式を含めたセル)を選択してください。
- リスクシミュレーター (Risk Simulator) | ツール(Tools) | 竜巻分析(Tornado Analysis)**を選択してください。
- 先例を確認し、適切な名称を付け直してください(先例を短い名称に付け替える事で竜巻とスパイダーチャートが見やすくなります)。その後、OK をクリックしてください。

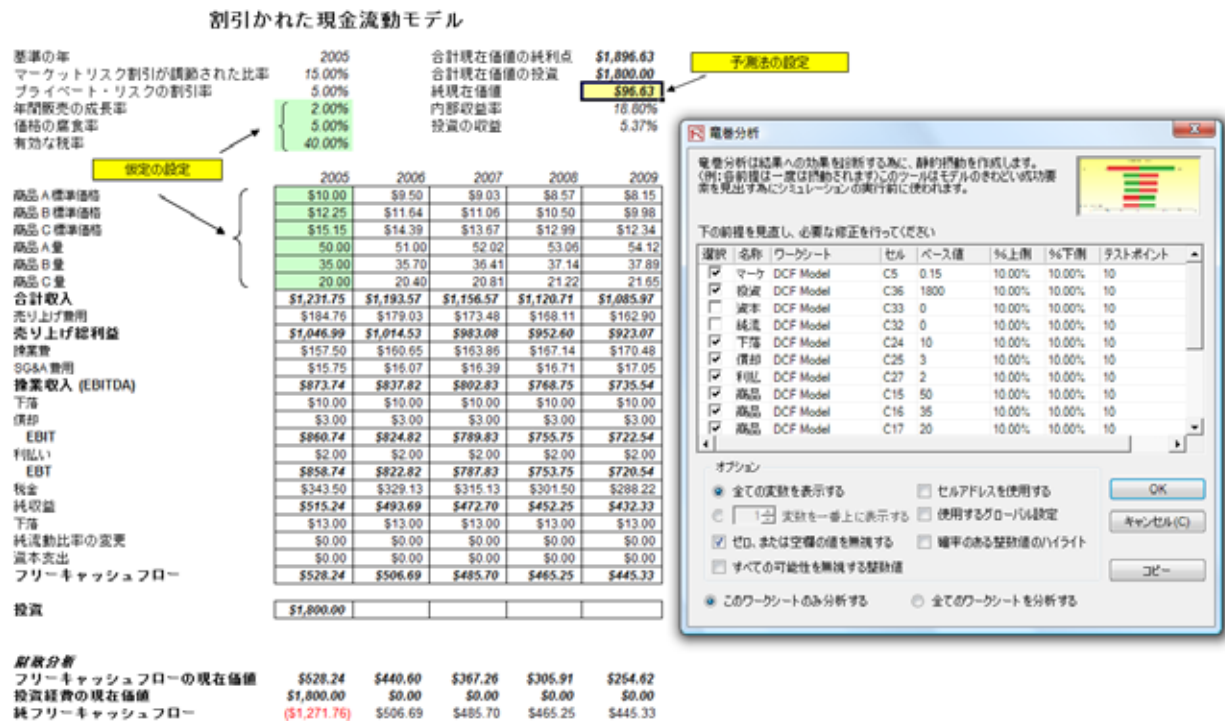


Figure 5.2 – 竜巻分析の実行

結果の解釈:

Figure 5.3 は、投資の資産が正味現在価値に影響を与える事を示し、税率、平均販売価格、要求された商品ラインの量などによって辿られた、竜巻分析の結果のレポートを表示します。レポートは4つの異なった要素を含んでいます。

- ① 実行された過程のリストの統計的な概略
- ② 感度表(Figure 5.4)は、初期の NPV の基本値とどうやって各入力 (例、投資は上回り+10%の\$1,800 から \$1,980 に変換され、下回りに-10% の\$1,800 から \$1,620 に変換) が変換されたかを表示しています。結果として成った NPV の上回り、そして下回りは、-\$83.37 と \$276.63 で、\$360 の総合変換で、NPV 上の最も高い影響を与える変数となります。先例変数は高い影響を与えるものからそうでないものの順でランクされます。
- ③ スパイダーチャート(Figure 5.5)は、これらの効果のグラフを表示しています。Y-軸は NPV の目的値であるのに対して、x-軸は、各先例の値(中心ポイントは、0%の変換で 96.63 の基本的なケース値で各先例の基本的な値です)上の百分位数の変換を表示しています。正の傾きの線は、正の関係、および効果を示し、負の傾きの線は、負の関係を示しています(例、投資は負としてスロープされ、投資のレベルが高いほど、NPV は低くなります)。スロープの絶対値は効果の大きさ(ステップラインは NPV 上での先例の x-軸の変換で与えられた y-軸での高い影響を示しています)を示しています。
- ④ 竜巻チャートは、このグラフを他の方法で表示し、最も高い影響を与えるものが最初に表示されています。x-軸は、NPV の値と基本的なケースの条件のチャートの中心です。チャートの緑の線は、正の効果を示しているのに対して、赤い線は、負の効果を示しています。したがって、投資の為に右側の赤線は高い NPV 上での負の投資を示しています。つまり、資本投資と NPV は、負の関係で関連されているという事です。逆に、価格と商品 A から C (チャートの右側にこれらの緑の線があります) の量には対立的な結果が見られます。

統計の概略

中でも竜巻チャートは有効なシミュレーションツールとされ、モデルの結果上での各変数の統計的なインパクトを獲得してくれます。すなわち、ツールは結果からなった混乱を最大から最少にランク付けし、モデルの予測または最終的な結果の変動を捕獲し、ユーザー指定の事前調整量モデルの各先例の変数を混乱等自動で行います。先例とはモデルの結果に影響を与える全ての記入と中間変数のことを示しています。例えば、モデルを $A=B+C$ とした場合、 $D=E+F$ で、 D と E が A の先例とされます（ C は先例ではなく計算からなった中間値であり）。混乱された値の範囲と数値はユーザーの指定で、予期された数値のまわり極値より小さい混乱でテストを設定することができます。現実的な環境では、極値は大きい、小さい、または不均衡なインパクトでははれられませんが、非線形は増加、または減少した経済のスケール、および規模のグループが変数の大小値のために起こる所に生じるとされています。そして、幅広い範囲だけがこの非線形のインパクトを捕獲することができるとされています。

竜巻チャートはモデルを先導する為、結果に最も影響を与える入力変数をはじめ全ての入力を表示します。チャートの標準は、同時にあり一貫した範囲の中（例、基準ケースから±10%、各先例の入力を混乱させ、基準ケースへのこれらの結果を比較することで表示されます。スパイダーチャートは一つの関係から幾つもの足が出ているようなことからスパイダー（クモ）を連想しています。肯定的な斜率ラインは肯定的な関係を示し、一方、否定的な斜率ラインは否定的な関係を示しています。それ以上に、スパイダーチャートは線形、または非線形の関係を表示するために使用できます。竜巻とスパイダーチャートは、シミュレーションする結果セル、および入力セルの重大な成功の要因の見出しを簡易してくれます。見つけ出された重大で不確実な変数の中からシミュレーションすべき変数が見つけ出されます。結果に於いて小さいインパクトを与える変数や、不確実に近い変数をシミュレーションする事で時間を無駄にしないで

結果

先例のセル	基準値: 96.6261638553219			入力変換		
	下側の出 力	上側の出 力	有効な範囲	下側の入 力	上側の入 力	ベースケ スの値
投資	276.62616	-83.37384	360.00	\$1,620.00	\$1,980.00	\$1,800.00
有効な税率	219.72693	-26.4746	246.20	96.00%	44.00%	40.00%
商品 A 標準価格	3.4265424	189.82679	186.40	\$9.00	\$11.00	\$10.00
商品 B 標準価格	16.706631	176.5457	159.84	\$11.03	\$13.48	\$12.25
商品 A 量	29.177498	170.07483	146.90	45.00	55.00	50.00
商品 B 量	30.533	162.71933	132.19	31.50	38.50	35.00
商品 C 標準価格	40.146587	153.10574	112.96	\$13.64	\$16.67	\$15.15
商品 C 量	48.047369	145.20496	97.16	18.00	22.00	20.00
価格の腐食率	116.80381	76.640952	40.16	4.50%	5.50%	5.00%
年間販売の成長率	90.588354	102.68541	12.10	1.80%	2.20%	2.00%
下落	95.084173	90.168155	3.08	\$9.00	\$11.00	\$10.00
利払い	97.088761	96.163566	0.93	\$1.80	\$2.20	\$2.00
償却	96.163566	97.088761	0.93	\$2.70	\$3.30	\$3.00
資本支出	96.626164	96.626164	0.00	\$0.00	\$0.00	\$0.00
経流動比率の変更	96.626164	96.626164	0.00	\$0.00	\$0.00	\$0.00

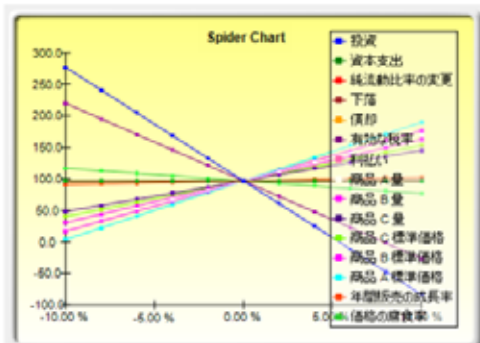


Figure 5.3 – 竜巻分析のレポート

メモ:

竜巻分析は静的な感度分析でモデルの各入力変数上に適用されるという事を忘れないで下さい。これは、各変数が個別的に攪乱され、結果としてなった効果は表に表示されます。これは、竜巻分析がシミュレーションを実行する前の鍵の要素とします。リスク分析の最初のステップとして、モデルのどこに最も重要な影響が得られ検出されるかを見出す事です。次のステップは、これらの重要な影響からどれが不確実かを見出す事です。この不確実な影響は、プロジェクトのクリティカルな成功のドライバーであり、モデルの結果はこれらのクリティカルな成功はドライバーに依存しています。これらの変数は、シミュレーションするべき変数と定められます。結果に小さな影響を与え、不確率だと定められる変数のシミュレーションに時間を取らないで下さい。竜巻チャートは、これらのクリティカルな成功ドライバーを早くそして簡単に見出す援助をします。この例証を辿ると、もし、要求された投資と有効税率の両方が事前に知られており、無変化だと仮定すれば、価格と量はシミュレーションされなければいけません。

先例のセル	基準価値: 96.6261638553219			入力変換		
	下側の出力	上側の出力	有効な範囲	下側の入力	上側の入力	ベースケースの値
投資	276.62616	-83.37384	360.00	\$1,620.00	\$1,980.00	\$1,800.00
有効な税率	219.72693	-26.4746	246.20	36.00%	44.00%	40.00%
商品 A 標準価格	34.255424	189.82679	186.40	\$9.00	\$11.00	\$10.00
商品 B 標準価格	16.706631	176.5457	159.84	\$11.03	\$13.48	\$12.25
商品 A 量	23.177498	170.07483	146.90	45.00	55.00	50.00
商品 B 量	30.533	162.71933	132.19	31.50	38.50	35.00
商品 C 標準価格	40.146587	153.10574	112.96	\$13.64	\$16.67	\$15.15
商品 C 量	48.047369	145.20496	97.16	18.00	22.00	20.00
価格の腐食率	116.80381	76.640952	40.16	4.50%	5.50%	5.00%
年間販売の成長率	90.588354	102.68541	12.10	1.80%	2.20%	2.00%
下落	95.084173	98.168155	3.08	\$9.00	\$11.00	\$10.00
利払い	97.088761	96.163566	0.93	\$1.80	\$2.20	\$2.00
償却	96.163566	97.088761	0.93	\$2.70	\$3.30	\$3.00
資本支出	96.626164	96.626164	0.00	\$0.00	\$0.00	\$0.00
純流動比率の変更	96.626164	96.626164	0.00	\$0.00	\$0.00	\$0.00

Figure 5.4 – 感度表

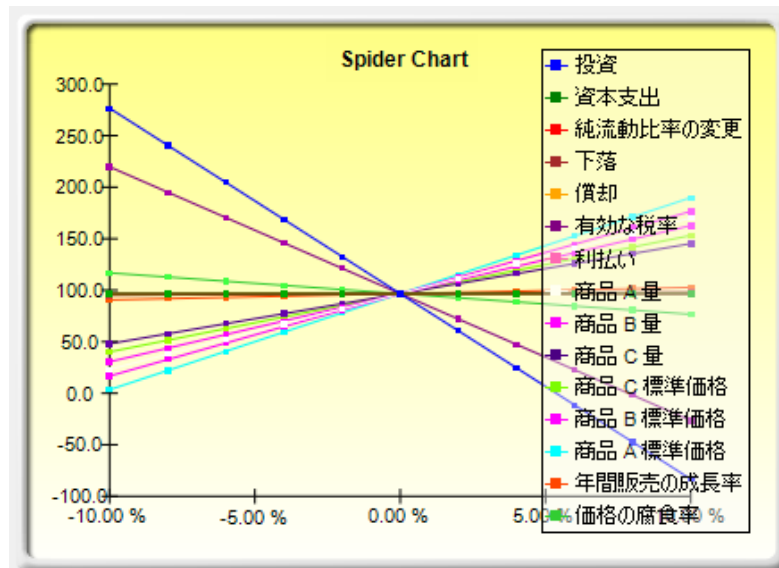


Figure 5.5 – スパイダーチャート

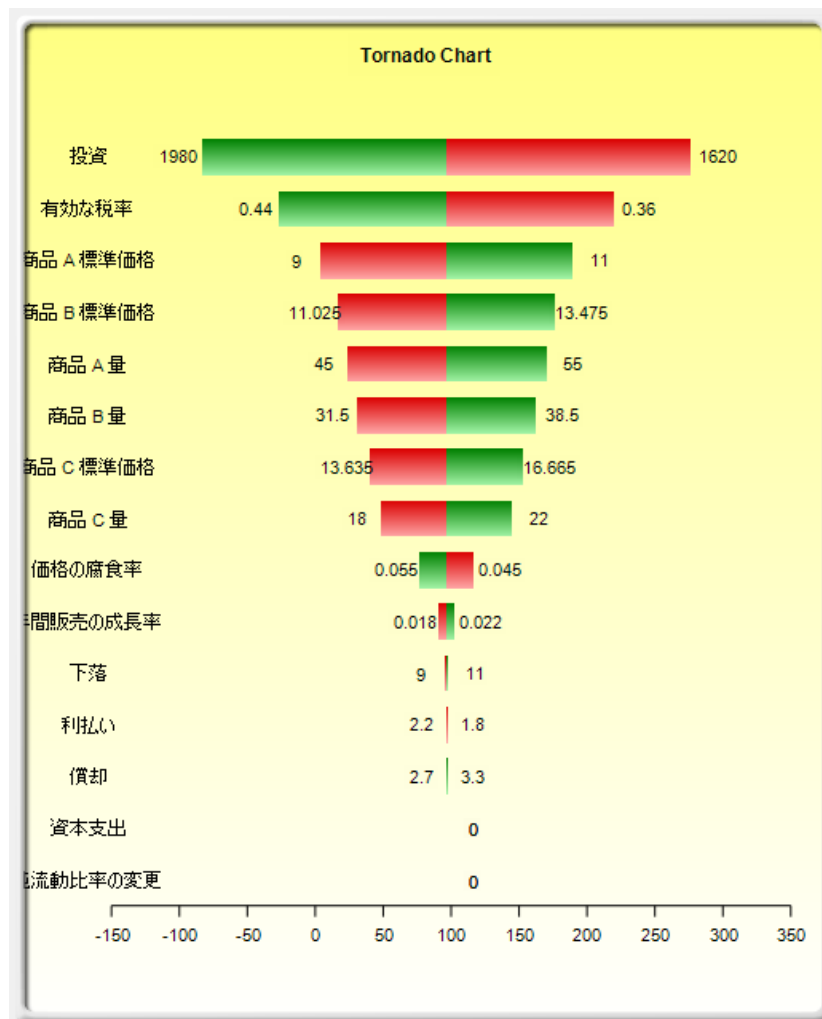


Figure 5.6 – 竜巻チャート

竜巻チャートは読解するのに簡単だとしても、スパイダーチャートでは、モデルに非線形があるかどうかを定める事が重要と成ります。例えば、Figure 5.7 は、非線形がかなり明白（グラフでの線は直線ではなく曲線です）な他のスパイダーチャートを表示しています。使用されたモデルは、**竜巻と感度チャート（非線形）**で、Black-Scholes オプションの価格モデルを例証モデルとして使用しています。非線形は、竜巻チャートから確認され、モデルにとって重要な情報となるか、モデルの活動の中で重要な判断メーカーを与えます。

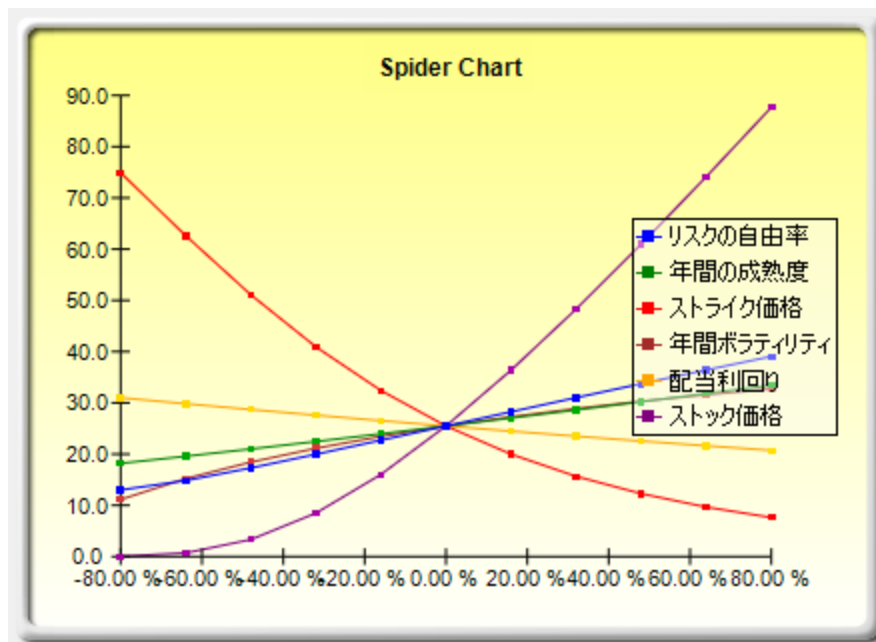


Figure 5.7 – 非線形的スパイダーチャート

感度分析

理論・セオリー:

関連する方法は感度分析です。竜巻分析(竜巻チャートとスパイダーチャート)がシミュレーションの実行前に静的な攪乱を適用するが他方で、感度分析は、シミュレーションの実行後に作成されたダイナミックな攪乱を適用します。竜巻とスパイダーチャートは、静的な攪乱の結果であり、各先例、または仮定変数は初期の量を一度に攪乱させ、変動の結果は表に表示されます。一方、感度チャートはダイナミックな攪乱の結果であり、複数の仮定は同時に攪乱され、モデルと変数の間の相関でのこれらの反復は、結果の変動で取得される事を意味しています。したがって竜巻チャートは、シミュレーションに

とってどの変数が最も適切で結果に導いてくれるかを見出してくれます。感度チャートは結果への影響を見出す一方、相互している複数の変数は、モデルで同時にシミュレーションされます。この効果は明確に Figure 5.8 で表示されています。クリティカルな成功のドライバーの順位表は、前に記述された竜巻チャートの例証に類似している事に注目してください。但し、仮定の間で相関が追加された場合、Figure 5.9 のように、とても異なった結果を表示します。例証に注目してください。NPV 上で価格の減少は小さな影響を及ぼしていますが、入力仮定のどれかが相関した場合、これらの相関する変数の間に存在する反復は、価格の減少がもっと大きい影響を及ぼすようにしてくれます。

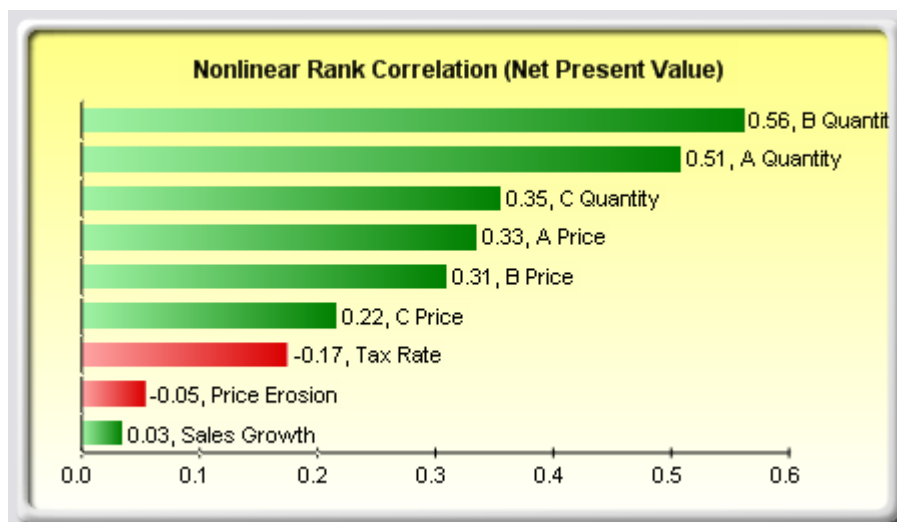


Figure 5.8 – 相関なしの感度チャート

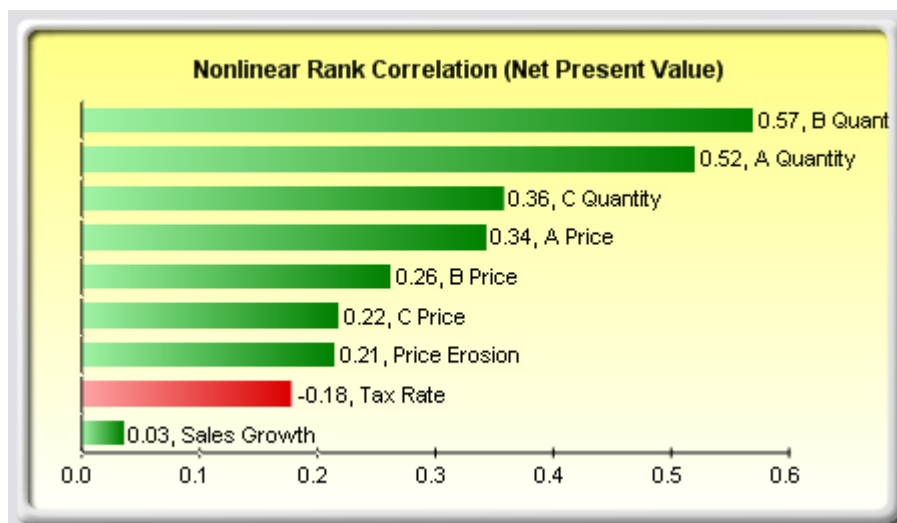


Figure 5.9 – 相関と感度チャート

手順

❏ モデルを作成するか開いてください。仮定と予測を定め、シミュレーションを実行してください (ここでの例証は、竜巻と感度チャート (線形) ファイルを使用しています)。

❏ リスクシミュレーター (Risk Simulator) | ツール (Tools) | 感度分析 (Sensitivity Analysis) を選択してください。

❏ 分析の為に予測を選択し、OK (Figure 5.10)をクリックしてください。

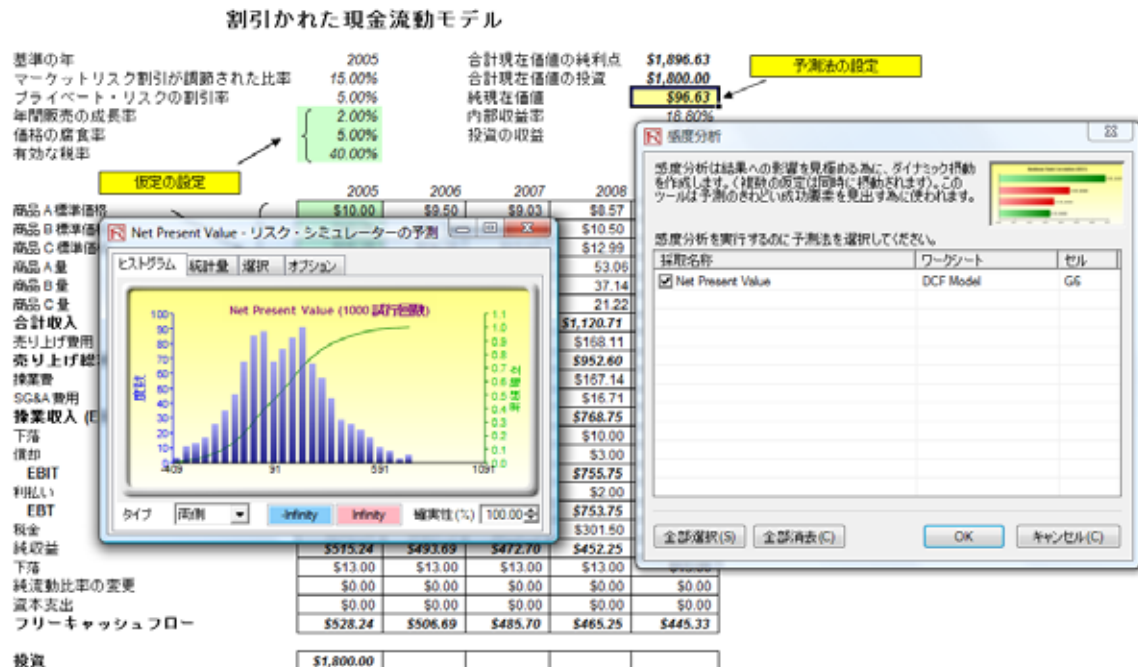


Figure 5.10 – 感度分析の実行

結果の解釈

感度分析の結果はレポートと二つのキーチャートを構成します。最初は、仮定予測の相関ペアを高いものから低いものへと、非線形順位表相関チャート (Figure 5.11)を構成します。この相関は非線形且つノンパラメトリックで、どの分布要件からも自由にします (例、ワイブル分布と仮定はベータ分布で比較できます)。このチャートからの結果は前に記述された竜巻分析に明白的に類似していますが (もちろん、知られている値だと初期に定められ、シミュレーションされなかった資本投資を除き)、例外が一つあります。感度チャート (Figure 5.11)では、税率は竜巻チャート (Figure 5.6)に比べて、最も低い位置に移行されました。これは、税率が有意な影響を及ぼすけれど、一旦、モデルでの他の変数が相互作用した後、税率の影響を減少させるからです (これは、税率は、履歴的税率の傾向には余り変動が無い様に、分布が小さく、また、税率は、税金前の収入の百

分位数の値がストレートで、他の先例の変数の方が大きな効果を持っているからです)。この例証はシミュレーション実行後の感度分析の適用は、モデル内でどのような相互関係があるか、そして特定の変数の影響が維持されているかを確認するのに重要であることを証明します。二つ目のチャート (Figure 5.12)は、説明された百分位数の変動を表示しています。すなわち、予測の変動の中で、どれほどの変動が、変数の間の全ての相互作用を考慮した各仮定によって説明されるかを示しています。説明された変動の足し算は常に 100% (時々、モデルに影響する他の要素がありますが。直接ここでは、採取できません)に近い事に注目してください。そして相関が存在するかどうか、そしてこれらの足し算は、時によって 100% (蓄積的な反復効果の実行の為)を超えます。

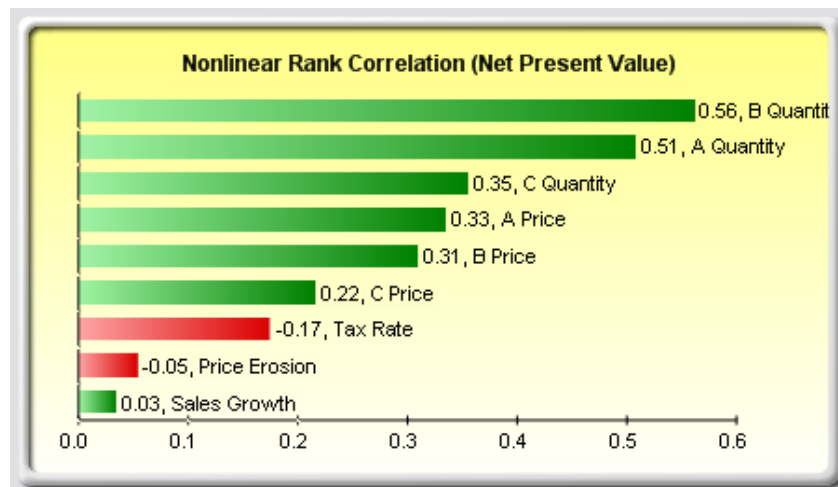


Figure 5.11 – 相関順位チャート

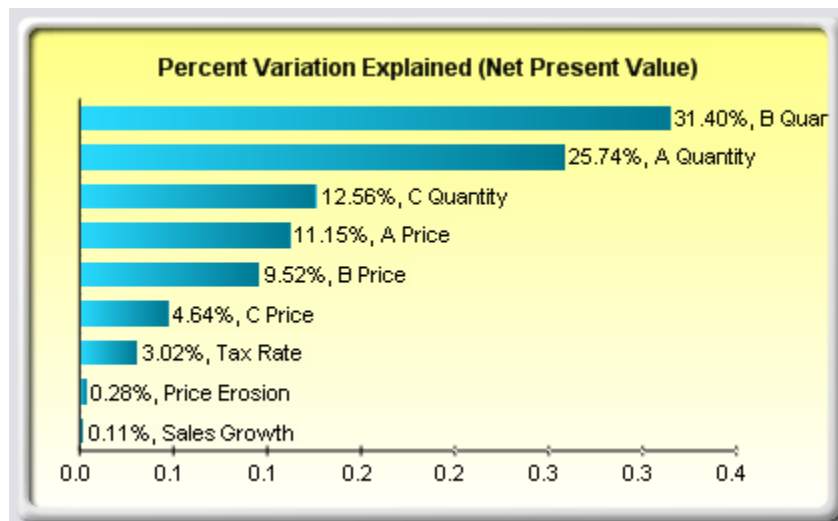


Figure 5.12 – 分散チャートへの貢献

メモ:

竜巻分析は、シミュレーションの実行前に遂行され、感度チャートはシミュレーションの実行後に遂行されます。竜巻分析でのスパイダーチャートは非線形を考慮する事が出来、感度チャートのランク相関は非線形と分布的に自由な条件が考慮できます。

分布的な適合: 単一変数と複数の変数

理論・セオリー:

他の強力的なシミュレーションツールは分布的な適合です。すなわち、どの分布が、モデルの特定の入力変数の分析の適用に当てはまるのでしょうか? どれが重要な分布的なパラメーターなのでしょうか? 履歴データが存在しない場合、分析家は疑問となっている変数の仮定を作成しなければいけません。一つの方法として、専門家のグループが各変数の振る舞いの推定を任せられる Delphi 方法を使用する事です。例えば、機械工学者のグループが検討、および厳格な実験を通してスプリングコイルの直径の極端な確率の評価を任せられるなど。これらの値は変数の入力パラメーター（例、正規分布と 0.5 と 1.2 の間の極値）として使用できます。検定が可能でない時（例、マーケットシェアと収入の成長率）、可能性のある結果の推定が管理によって作成でき、最適なケース、最もおおよそなケース、そして悪いケースのシナリオを与えてくれます。

但し、信頼できる履歴データが存在する時は、分布的な適合を行うことができます。履歴パターンと履歴傾向自身の反復は、最適な適合分布と、シミュレーションされる変数の最適な定義の重要なパラメーターの検出に使用されます。Figures 5.13 から 5.15 は、分布的な適合の例証を表示しています。このイラストは、例証ファイルの **データの適合** ファイルを使用しています。

手順:

- 分布シートを開き、適合させたいデータを呼び出してください。
- 適合させたいデータを選択してください（データは一つの縦列に複数のセルに記入されてなければいけません。）
- **リスクシミュレーター (Risk Simulator) | ツール (Tools) | 分布的な適合(単一変数) (Distributional Fitting (Single-Variable))** を選択してください。
- 適合させたい、ある特定の分布を選択するか、全ての分布が選択されているデフォルトを維持し、OK (Figure 5.13)をクリックしてください。

- 適合の結果を確認し、希望する重要な分布を選択し、OK (Figure 5.14)をクリックして下さい。

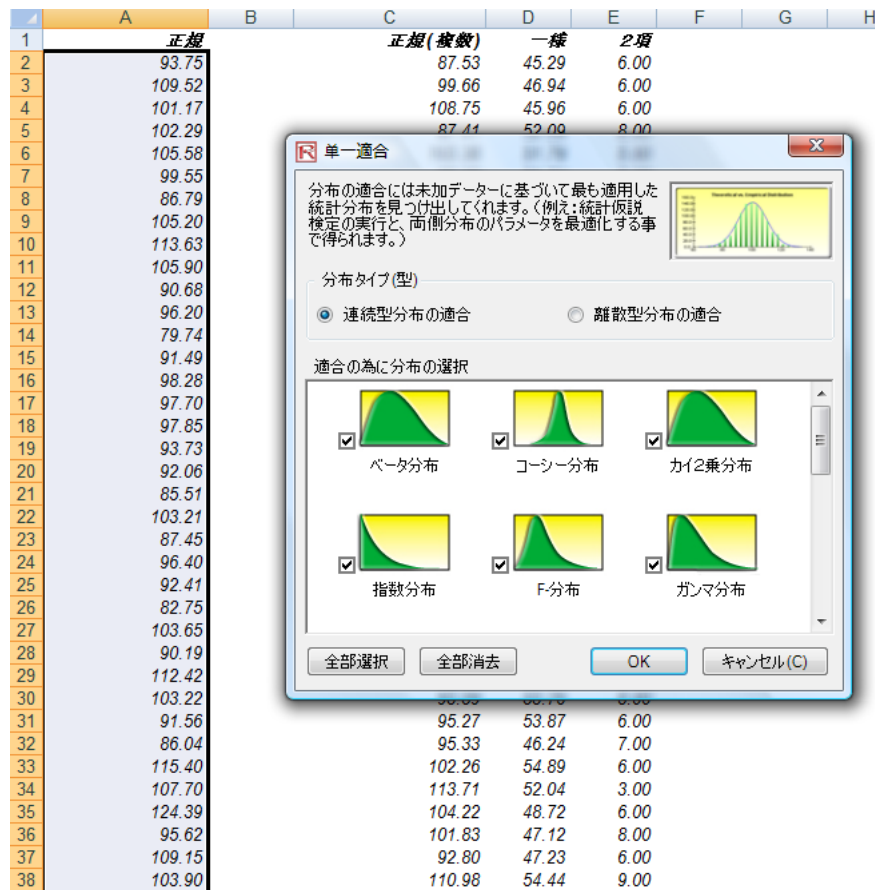


Figure 5.13 – 単一変数の分布的な適合

結果の解釈:

無帰仮説は、適合対象の分布が、適合対象のサンプルデータが引き出される母集団と同じ分布を持っているかどうかを検定されます。したがって、計算された p -値がクリティカルなアルファレベル(一般的に 0.10 か 0.05)よりも低い場合は、分布は間違っている事を示しています。一方、 p -値が高ければ最も良く分布はデータに適合します。たいてい、 p -値を説明された百分位数として考える事が出来、 p -値が 0.9727 (Figure 5.14) の場合、99.28 の平均値と 10.17 の標準偏差を持つ正規分布が、データの変動の 97.27%を説明した事になり、特別に良い適合を示します。どちらの結果(Figure 5.14)も、レポート(Figure 5.15)も検定統計、 p -値、論理的な統計 (選択された分布に基づき)、経験的な統計 (未加工データに基づき)、初期のデータ(使用されたデータの記録の維持)と完成された仮定と重要な分布のパラメーター(例、自動的な仮定の生成の選択とシミュレーション

プロファイルが既に存在する場合)を表示しています。また、結果は選択された全ての分布とどれほどデータに適切に適合しているかを順位表に示します。

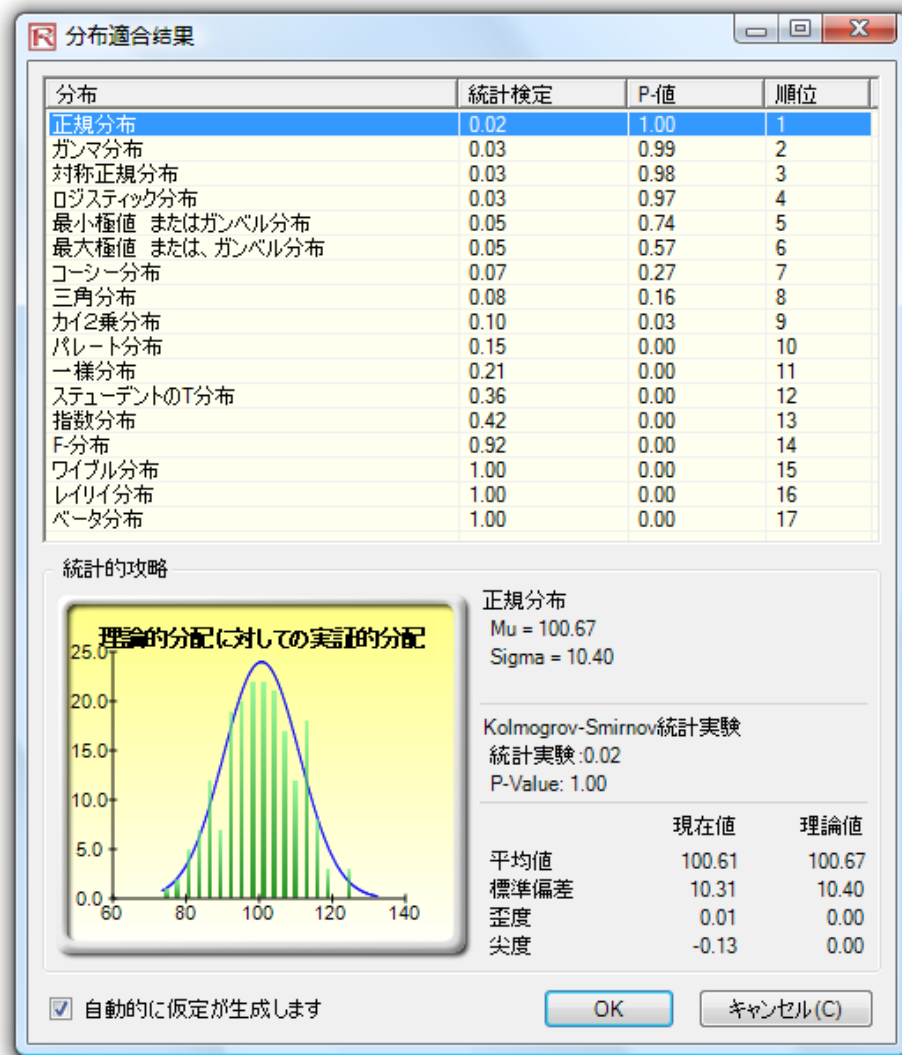


Figure 5.14: 分布的な適合の結果

統計の概略

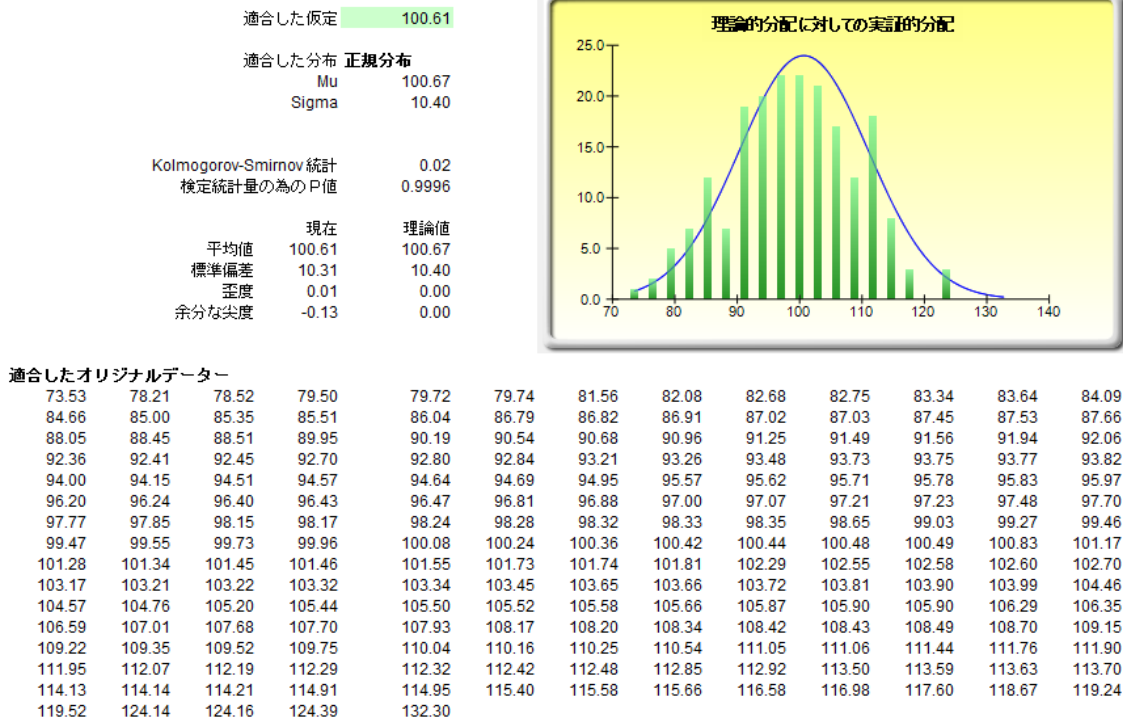


Figure 5.15: 分布的な適合のレポート

複数の変数の適合過程は、単一変数の適合の時のようにかなり似ています。いずれにしても、データは列に編集されていなければいけません（例：各変数は一つの列とする）すると、他の変数は一度に全て適合されます。

手順:

- 分散シートを開き、適合したいデータと呼び出してください。
- 適合したいデータを選択してください（データは複数の縦列と横列に入っていないければいけません。）
- リスク・シミュレーター (Risk Simulator) | ツール (Tools) | 分布的な適合 (複数の変数) (Distributional Fitting (Multiple Variables)) を選択してください。
- データを見直し、関連のある分布タイプを選択し、OK をクリックしてください。

メモ:

分布的な適合ルーチンに用いられている統計順位方法はカイ 2 乗検定と Kolmogorov-Smirnov 検定を使います。前者は離散系分布を後者は連続分布を検定する為に使います。簡潔にいうと、内部最適化ルーチンと合わせた仮説検定は検定された各分布の最適なパラメーターを見つけ出す為に使用します。そして、結果は順位表に記されます。

ブートストラップシミュレーション

理論・セオリー:

ブートストラップシミュレーションは予測統計の精度、信頼性やサンプル未加工データなどを推定する簡単なテクニック法です。基本的に、ブートストラップは仮説検定に使用されます。古典的な方法論ではサンプル統計の信頼性を記述する為に数学的な方式を使用していました。この方法はサンプル統計の分布が正規分布へ統計の標準エラー、または信頼区間の計算に割りと簡易に近づけてくれます。いずれにしろ、統計のサンプル分布が正規分布していない、または簡易に見つからない場合はこの古典的な方法は無効か、使用が困難になります。逆に、ブートストラップがサンプル統計データを何回も繰り返しサンプルし、それらのサンプルを基に様々な統計の分布を作り出し、実践的に分析をします。

手順:

- シミュレーションを実行してください。
- リスク・シミュレーター (*Risk Simulator*) | ツール (*Tools*) | ノンパラメトリック ブートストラップ (*Nonparametric Bootstrap*) を選択してください。
- ブートストラップの実行の為一つだけの予測を、統計を選択し、ブートストラップ試行数を記入し OK をクリックしてください。(参照 Figure 5.16)

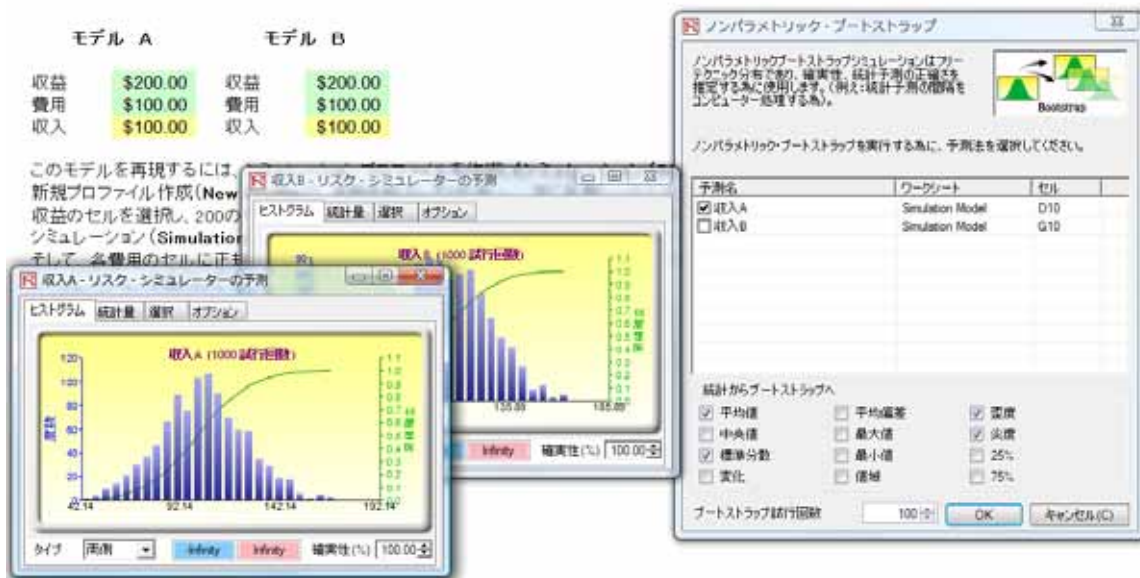


Figure 5.16 – ノンパラメトリック ブートストラップの結果

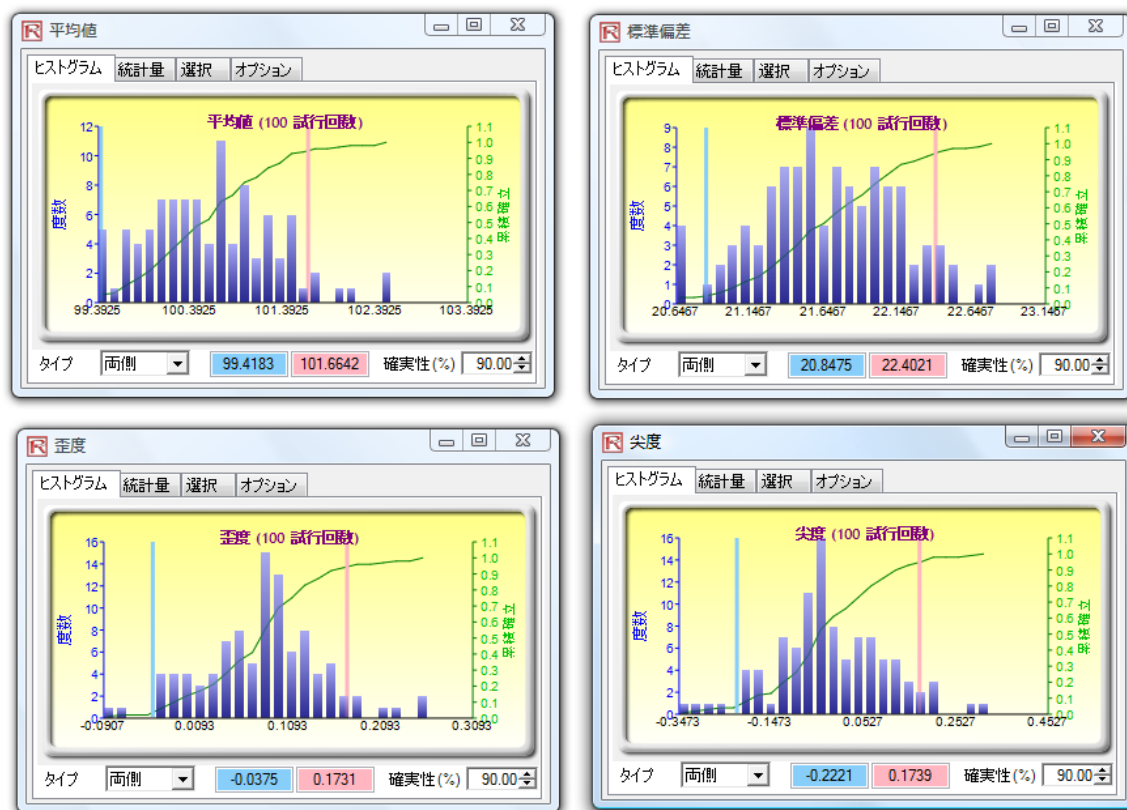


Figure 5.17 – ブートストラップシミュレーションの結果

結果の解釈:

基本的に、ノンパラメトリック ブートストラップはシミュレーションに基づいたシミュレーションと言えるでしょう。シミュレーションの実行後、結果から起こる統計が表示されます。しかし、各統計の精度と統計的な意味は時々解決されないままの場合があります。例えば、シミュレーション実行の歪度の統計が -0.10 の場合、本当に分布は負の歪なののでしょうか、またわずかな負の数はランダム試行に起因するのでしょうか？ちなみに統計が -0.15 , -0.20 の場合はどうなのでしょう？すなわち、どれくらい十分に遠くになれば、分布に負の歪があると考えられるのでしょうか？この疑問は他の統計にも適用できます。一つの分布は統計的に他の分布と同じとすれば統計的な関係が計算できるのでしょうか、またこれらはまったく相違的なのでしょうか？Figure 5.17 はブートストラップのサンプルを参照しています。また、90%の統計の歪度の信頼は -0.0189 から 0.0952 の間にあり、0 値も信頼度の中に入ります。つまり、90%の信頼は統計の歪度は統計的にゼロと違いがないことを示しており、この分布は対称的で歪がないと考えていいということです。逆に値が 0 を超えている場合、その反対を示し、分布は非対称的で歪んでいることを示します。（予測統計が正の場合は歪が無く、予測統計が負の場合は歪みも負の可能性もあります）。

メモ:

ブートストラップの由来は“自分の独力で自分自身を引き上げる”と言う意味があり、この方法は、統計学の精度を分析するのに統計自体の分布を使用します。ノンパラメトリックシミュレーションは各ゴルフボールが履歴データポイントに基づいているとして、交換で大きいかごからゴルフボールを単にランダム方式に選ぶことです。例えば、カゴの中に 365 個のゴルフボールが入っていると（365 の履歴データポイントを表しているとします）します。想像してみてください、ランダム方式で選び出された各ゴルフボールの値がホワイトボードに書かれるとします。交換されながら選ばれたゴルフボールの結果はボードの 365 数の列の一行目に記入し、重要な統計がこれらの 365 列上で（例：平均値、中間値、標準偏差など）計算されます。そして、例えば過程が 5,000 回繰り返されたとします。そうとなるとホワイトボードは 365 列と 5,000 の縦列に記入することになります。したがって、統計の結果は（これは 5,000 の平均値、5,000 の中間値、5,000 の標準偏差などがあるということを認識し）表に記され、それぞれの分布は表示されます。重要な統計は後ほど、表に表示されます。この結果によって、シミュレーションの予測の信頼度が確認できます。つまり、10,000 回試行のシミュレーションでは、結果として起こる予測平均が 5.00 ドルであるということを前提とすると、解析者はどの程度結果に対して精度を持つのでしょうか？ブートストラップは計算された統計平均値

を統計の分布の表示も含めての、信頼区間を確かめさせてくれます。最後に、ブートストラップの結果は、統計の大数の法則と中心極限定理によると、サンプル平均値の値はサンプルのサイズが増えると、不偏推定量で、真の母集団平均値にほぼ近づき、等しくなると分かります。

仮説実験

理論・セオリー:

仮説検定は二つの分布の平均値や分散の検定を行っているときに、それらが統計的に同じか、異なるかを示してくれます。まったく違う二つの予測の平均値や分散の違いはランダム試行か、実際に統計的にまったく違う値に基づいているからだと考えられます。

手順:

- シミュレーションを実行
- リスク・シミュレーター (*Risk Simulator*) | ツール (*Tools*) | 仮説実験 (*Hypothesis Testing*) を選択
- 同時に実験を実行するのに二つの予測を選択肢、仮説実験のタイプを選択し、OK をクリックしてください(Figure 5.18)。

	モデル A		モデル B
収益	\$200.00	収益	\$200.00
費用	\$100.00	費用	\$100.00
収入	\$100.00	収入	\$100.00

このモデルを再現するには、シミュレーションプロファイル、新規プロファイル作成 (**New Profile**), 名称を記入して、収益のセルを選択し、200の平均値を正規分布に与え、シミュレーション (**Simulation**) | 仮定の設定 (**Set Ass**)。そして、各費用のセルに正規仮定を定義してください。シミュレーションを実行してください。

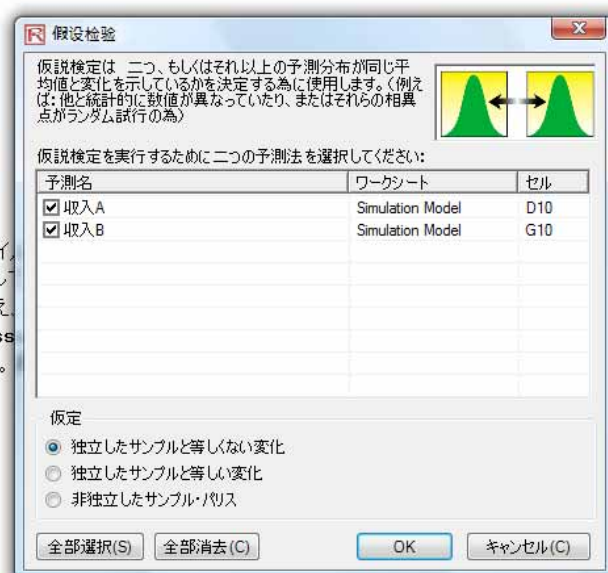


Figure 5.18 – 仮説実験

レポートの解釈:

両側仮説の検定は帰無仮説上で実行されています(H_0)。この二つの母集団(変数)は統計的には似ていると言う事を意味していて、対立仮説(H_a)は二つの母集団(変数)は総計的に似ていないと言うことを意味しています。もし、計算された p 値が 0.01, 0.05, または 0.10 と同じ、または小の場合 ($p \leq$) は、帰無仮定が拒絶されていることを示し、予測が総計的に 1%、5%、10%とはレベルが異なっていることが分かります。一方、もし帰無仮説が拒絶されてなく p 値が高い大のばあい ($p \geq$) は、二つの予測分布は総計的に似ていると言うことが分かります。同じような分析の実行は F・検査の pairwise を使用し、同時に二つの相違仮説の予測を行うことでできます。P 値が小の場合、分散（それから標準偏差）が総計的に異なっていて、 p 値が大の場合は分散が総計的に似ていると判断できます。

二つの予測の平均値と変化の仮説検定

2列に並んだ仮説検定は、二つの変数の母集団が統計的に相似している帰無仮説 H_0 で実行されます。また、対立仮説の場合は、

仮説検定は、平均値と変化の検定で、二つの分布が統計的に相似しているか、または、総計的に異なっているかを決定する為に使います。つまり、二つの平均値と変動を比較するという事です。つまり、二つの平均と二つの変動に起こる違いを見る為に、ランダムチャンスが実際にこれら互いに異なっている事に基いているということである。等しくない変動の2可変的なテストは(予測1の母分散は予測2の母分散と異なる予期した場合)予測の配分が異なった母集団から成るときに、適切である(例えば、2つの異なった地理的位置から採取したデータ、2つの異なった事業の作業など)。等しい変動の2可変的なテストは、(予測1の母分散は予測2の母分散と等しい予期した場合)予測の配分が同じような母集団から成るときに適切である(例えば、同じような指定を含む2つのエンジンの設計から採取したデータなど)。組み合わせられた、二つの独立可変的なテストは予測の配分が同じような母集団から成るときに、適切です(例えば、顧客の同じグループから異なった機会に採取したデータなど)。

2列に並んだ仮説検定は、二つの変数の母集団が統計的に相似している帰無仮説 H_0 で実行されます。また、対立仮説の場合は、母集団が統計的に異なっている事を示します。計算された p -値が 0.01, 0.05, 0.10より小、または同じの場合、仮説が拒絶された事を意味します。つまり、予測の平均値は統計的に、1%5%と10%の有意水準に対して、かなり異なっているということです。帰無仮説が拒絶されていない場合は、 p 値が高く、二つの予測分布の平均値は、統計的に他と似ています。同じ分析は、二つの予測の変動の際に、pairwise F検定を使って実行されます。 p 値が小さく、変動(と、標準偏差)が統計的に他と異なっているとすれば、 p -値が大きい場合は、変動が統計的に他と似ている事を示します。

結果

仮説検定の仮定	不平等な変化
t-統計の計算:	1.3842404
t-統計の為のP-値:	0.1664397
F-統計の計算:	1.0896575
F-統計の為のP-値:	0.1749742

Figure 5.19 – 仮説検定の結果

メモ:

二つの変数の t 検定と等しくない分散（予測 1 の母集団分散は予測 2 の母集団分散と異なると予期し）は、予測の母集団分散が異なるときに最適です（二つの異なった現場から収集したデータは、二つの異なったビジネス営業ユニットとフォースが必要）。二つの変数の t 検定と等しい分散（予測 1 と予測 2 の母集団分散は同じと予期し）は、予測分布が同じ母集団の場合に最適です（二つの異なったエンジン・デザイン（似た明細を加え）から収集したデータ）。二つのペア化され、従属された t 検定は予測分布がまっ

たく同じ母集団の場合に最適です（同じ現場、異なったシチュエーションで同じグループのお客さんから収集したデータ）。

データ採取とシミュレーション結果の保存

リスク・シミュレーターのデータ採取の手順を使用すれば容易にシミュレーションの未加工データの採取ができます。仮定のデータも予測のデータも採取できますが、あらかじめシミュレーターが実行されていることを確認してください。採取されたデータは後ほど、多様の他の分析に使用できます。

手順:

- 模型を開くか作成した後、仮定と予測を設定し、シミュレーションを実行してください。
- **リスク・シミュレーター (Risk Simulator) | ツール (Tools) | データ採取 (Data Extraction)** を選択してください。
- データの採取したい仮定、予測、または両方を選択し、OK をクリックしてください。

採取されたデータ様々なフォーマットで保存できます。

- 未加工データで新しいワークシートに表示され、シミュレートされた値（仮定そして予測の両方）は、保存、または分析ができます。
- フラット テキスト ファイル (flat text file)。データは他の分析ソフトウェアにエクスポートされます。
- リスク・シミュレーター ファイル。 **リスク・シミュレーター (Risk Simulator) | ツール (Tools) | データを開く・インポート (Data Open/Import)** を選択（仮定でも予測でも）すれば いつでも結果を回復することができます。

三つ目のオプションは、シミュレートされたデータ結果を*.risksim ファイルとして保存することです。データ結果はいつでも回復することができますし、シミュレーションを毎回実行させる必要がありません。Figure 5.21 のダイアログ・ボックスではデータの採取、エクスポート、採取データ結果の保存の画面が表示されています。

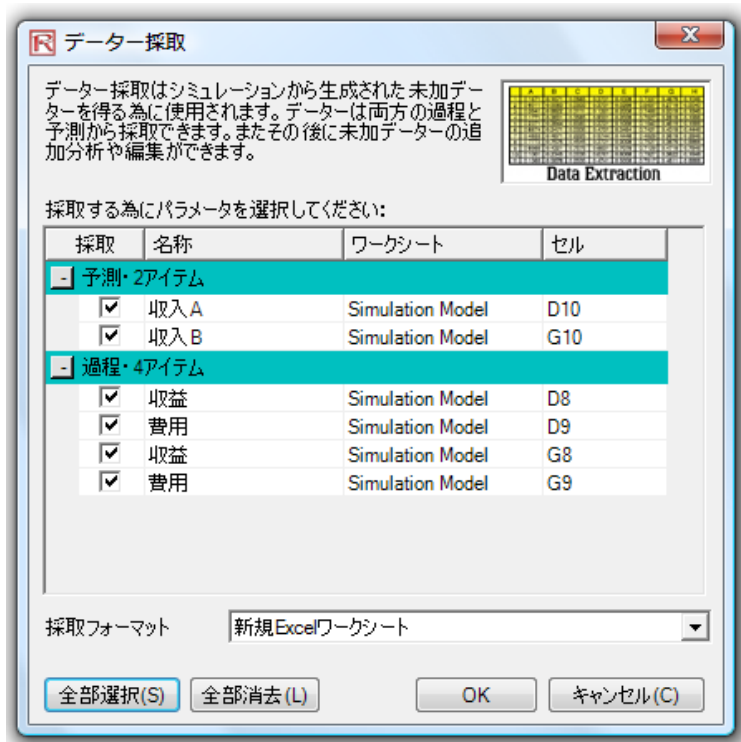


Figure 5.21 – シミュレーションレポートのサンプル

レポートの作成

シミュレーション実行後、設定した仮定、予測、シミュレーションの結果等のレポートの作成ができます。

レポートの作成の仕方:

- モデルを開くか作成した後、仮定と予測を設定し、シミュレーションを実行してください。
- **リスク・シミュレーター (Risk Simulator) | レポートの作成 (Create Report)** を選択してください

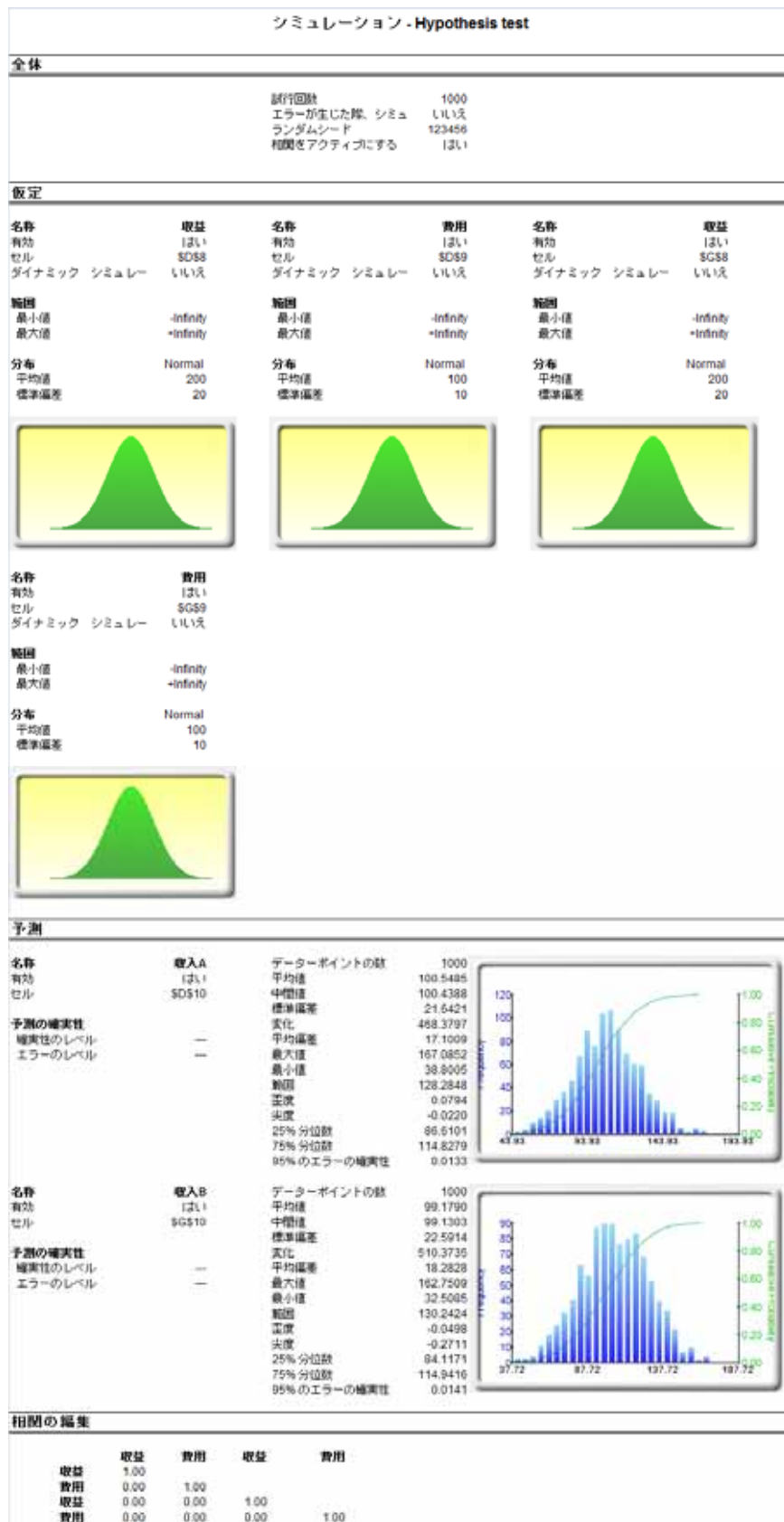


Figure 5.21 – シミュレーションレポートのサンプル

リスクシミュレーターでのこの高度な分析ツールは、データの計量経済学の性質を定める為に使用されます。診断は、不均一分散、異常値、指定されたエラー、小数、季節性と確率の性質、正常性、エラーの球形と多重共線性を含みます。各検定は、それらのモデルのレポートで細かく記述されています。

手順

- 例証モデルを開き(リスクシミュレーター (Risk Simulator) | 例証 (Examples) | 回帰診断 (Regression Diagnostics))、時系列データのワークシートで変数名称を含んだデータを選択してください(セル C5:H55)。
- リスクシミュレーター (Risk Simulator) | ツール (Tools) | 診断ツール (Diagnostic Tool) を選択してください。
- データを確認し、下記のメニューから従属変数 Y を選択して、終了後、OK をクリックして下さい(Figure 5.22)。

重回帰分析データー設定

従属変数 Y	変数 X1	変数 X2	変数 X3	変数 X4	変数 X5
521	18308	185	4.041	79.6	7.2
367	1148	600	0.55	1	8.5
443	18068	372	3.665	32.3	5.7
365	7729	142	2.351	45.1	7.3
614	100484	432	29.76	190.8	7.5
385	16728	290	3.294	31.8	5
286	14630				
397	4008				
764	38927				
427	22322				
153	3711				
231	3136				
524	50508				
328	28886				
240	16996				
286	13035				
285	12973				
569	16309				
96	5227				
498	19235				
481	44487				
468	44213				
177	23619				
198	9106				
458	24917				
108	3872				
246	8945				
291	2373				
68	7128				

診断ツール

多変数の設定によって起こる予測問題に対する診断を行うツールです。

従属変数 従属変数 Y

従属変数 Y	変数 X1	変数 X2	変数 X3	変数 X4	変数 X5
521	18308	185	4.041	79.6	7.2
367	1148	600	0.55	1	8.5
443	18068	372	3.665	32.3	5.7
365	7729	142	2.351	45.1	7.3
614	100484	432	29.76	190.8	7.5
385	16728	290	3.294	31.8	5
286	14630	346	3.287	678.4	6.7
397	4008	328	0.666	340.8	6.2
764	38927	354	12.938	239.6	7.3
427	22322	266	6.478	111.9	5
153	3711	320	1.108	172.5	2.8

OK

キャンセル(C)

Figure 5.22 – データの診断ツールの実行

予測と回帰分析の一般的な違反は不均一分散で、エラーの変動は時間と共に増加することを示しています。(参照、Figure 5.23 診断ツールを使用した検定結果の表示)。一見、縦軸のデータの分散の幅は増加、および時間と共に超えてしまい、一般的に、決定係数 (R- 平方された係数) は不均一分散が存在する時に有為に下降します。従属変数の分散が一定していない時、エラーの分散も一定しません。従属変数の不均一分散が強調するとこれらの効果はそれほど深刻ではなく：最小二乗推定はまだバイアスされてなく、傾きと切片の推定は、エラーが正規分布されている場合は、どちらも正規分布され、また、エラーが正規分布されていない場合、小さい正常漸近的な分布（データポイントの数が大きくなるように）が見られます。傾きと全体的の変数の値の分散の推定は、不確実な可能性があります、従属変数の値がこれらの平均値に対照的であれば、不確実性はあまり大きくなる可能性はありません。

もしもデータポイントが小さい(小数)場合は、仮定の違反の検出が困難になります。小さいサンプルサイズでは、非正規性、および分散の不均一分散のような仮定の違反が存在するとしても検出が難しくなります。小さい数のデータポイントと線形回帰は過程の再度の違反からの防御は余り強くありません。少ないデータポイントだと、データに適合する線がマッチするか、また、どの非線形公式が最も適切かを定めるのが難しくなります。一つもの仮定が違反されなかったとしても、データポイントの小数上の線形回帰は、傾きとゼロの間の、例え、傾きがゼロに同一していないとしても、違いを見分ける十分な力を持っていません。この力は残ったエラー、独立変数で観測された分散、検定のアルファレベルの選択された有意性とデータポイントの数によって変わります。力の減少は、残余の分散の増加、有意レベルの減少（例、検定が最も厳しくされるように）のように減少し、観測された独立変数の増加での分散の増加とデータポイントの数の増加のような増加のように増加します。

値は、異常値の存在がある為、同様に分布していないかもしれません。異常値とはデータの非正常な値を示します。異常値は、適合された傾きと切片上に強力な影響を及ぼし、データポイントの大きさへ乏しい適合を与えます。異常値は、残余の分散の推定を増加する傾向を見せ、帰無仮説の拒絶のチャンスを低くします。例えば、予言エラーの確率を高くするなど。エラーの記録は、修正が可能でなければいけなく、また、同じ母集団からサンプルされなかった従属変数の値の記録は重要です。一見、異常値は、同じ母集団からではあるが非正常な母集団からの従属変数の値に由来するかもしれません。但し、

ポイントとして、分散プロットでの異常値の必要はない独立、および従属変数のどちらも珍しい値となります。回帰分析では、適合されたラインは異常値に非常に敏感で無ければいけません。つまり、少ない平方の回帰は異常値に対して抵抗が無く、したがって、どれも適合した傾きの推定となりません。他のポイントから縦軸に削除されたポイントは、データの残余の全体的な線のトレンドを辿る代わりとしても、特にポイントがデータの中心から横線的に離れている場合、適合された線をポイントの近くに通す原因となります。

但し、異常値が削除された場合、良い方法を選ぶ必要があります。しかし、ほとんどのケースで異常値が削除された場合、回帰の結果は向上しますが、まず経験的な説明が存在しなければいけません。例えば、ある特定の会社の株式リターンの実行を戻すと、株式市場での下降による異常値は含まれていなければいけません。これらは、ビジネスサイクル内ではどうしようもないような本当の異常値ではありません。これらの異常値を見合わせ、回帰方式を使用して会社の株式に基づいた一人の退職基金を予測するのは、最大な不正確な結果を齎します。一方、異常値は、不正常的なビジネスコンディション（例、合併および獲得）によって生じたとし、このビジネスの構成の変換は、繰り返さないように予測されていない為、これらの異常値は削除する事が出来、最初にデータを清潔にしてから回帰分析を実行します。ここでの分析は、異常値しか検出できず、ユーザーの好みによってこれらを残すかそれとも排除するか選択できます。

時々、従属、そして独立変数の非線形的な関係は、線形的な関係よりも適切です。どのケースにしても、線形的な回帰を実行する事は最適ではありません。もし線形モデルが正しくなければ、傾きと切片の推定と線形的な回帰から適合された値はバイアスされ、適合された傾きと切片の推定は有為的ではありません。独立、または従属変数の限られた範囲を超え、非線形モデルは、線形モデルによって近似（線形的な補入のバイアスの要素）されますが、精度のある予言の為には、データにとって適切なモデルの選択が必要となります。回帰の実行前に、データに非線形の変換がまず適用されなければいけません。独立変数（他の方法は、平方根、および独立変数を二つ目、または三つ目の力に上昇する）の自然対数を取る事で、回帰、または予測を非線形的に変換されたデータを使用して実行する事です。

変数 Y	不均一分散		小多数性		アウトライヤー		非線形性	
	W-検定 p-値	仮説検定 結果	近似 結果	自然 下限	自然 上限	可能なアウトライヤー の数	非線形検定 p-値	仮説検定 結果
変数 X1	0.2543	Homoskedastic	問題なし	-21377.95	64713.03	2	0.2458	線形
変数 X2	0.3371	Homoskedastic	問題なし	77.47	445.93	2	0.0335	非線形
変数 X3	0.3649	Homoskedastic	問題なし	-5.77	15.69	3	0.0305	非線形
変数 X4	0.3066	Homoskedastic	問題なし	-295.96	628.21	4	0.9298	線形
変数 X5	0.2495	Homoskedastic	問題なし	3.35	9.38	3	0.2727	線形

Figure 5.23 – 異常値の検定結果、不均一分散、小数と非線形

時系列データを予測する為のほかの一般的な方法は、独立変数の値が本当に各自独立しているか、または従属しているかどうかを確認する事です。時系列を通して採取された従属変数の値は自己相関されなければいけません。連続的に相関された従属変数の値、傾きと切片の推定はアンバイアスされますが、これらの予測の推定と分散は不確実な為、特定の統計的な最良な適合の検定の評価は誤りとなります。例えば、利率、インフレーションの比率、販売、収入と時系列データに関連する他の比率等、現在の周期の値は先周期の値に関連され（明らかに三月のインフレーションは 2 月のレベルに関連され、これ自体が一月のレベルに関連されるなど）、一般的に自己相関されます。この関係を見無視する事は、バイアスをそして精度の少ない予測を齎す事になります。どのイベントでも、自己相関回帰モデル、または ARIMA モデルはより良く適用 (リスクシミュレーター (Risk Simulator) | 予測(Forecasting) | ARIMA)されなければいけません。最後に、シリーズの自己相関の関数は非季節性で、滑らかに下降する傾向を持っています (モデルで非季節性レポートを参照してください)。

もし、自己相関 $AC(1)$ がゼロと等しくない場合、これはシリーズがまず最初に連続的に相関されている事を示しています。もし、 $AC(k)$ が、増加する時間差と幾何学的に少し離れている場合、シリーズは低順位の自己回帰過程を辿っている事を示します。もし、 $AC(k)$ が、小さい数の時間差の後にゼロに下降した場合、シリーズは低順位の移動平均過程を辿っている事を示しています。一部分的な相関 $PAC(k)$ は、介入した時間差から相関を削除した後の k 周期の値の相関を測定します。自己相関のパターンは、 k より少ない順位の自己回帰によって得る事が出来、時間差 k への一部の自己相関は、ゼロに近似するはずで、Ljung-Box Q-統計と時間差 k に対するこれらの p -値は、帰無仮説を持ち、 k 順序を上回る自己相関はありません。自己相関のプロットの点線は、近似的な二

つの標準エラーの範囲です。もし自己相関がこれらの範囲を持っていない場合、ゼロから有為的な異なりが無く、5%の有意性レベルを示しています。

自己相関は、従属している Y 変数自身の過去への関係を測定します。分布的な時間差は逆に、従属している Y 変数と異なった独立している X 変数の間の時間差の関係です。例えば、死亡率の分散と方向は、時間差(一般的に 1 から 3 ヶ月)に対してフェデラルファーズの比率を辿る傾向見せます。時により時間差はサイクルと季節性(例、アイスクリームの販売は夏に増加する傾向を見せる為、過去 12 ヶ月の最後の夏の販売に関連されます)を辿ります。分散された時間差分析(Figure 5.24)は、様々な時間差に対してどのように従属変数が各独立変数の関連されているか、どの時間差が統計的に有意で考慮されるのかを定める為に、いつ全ての時間差が同時に定められるのかを表示しています。

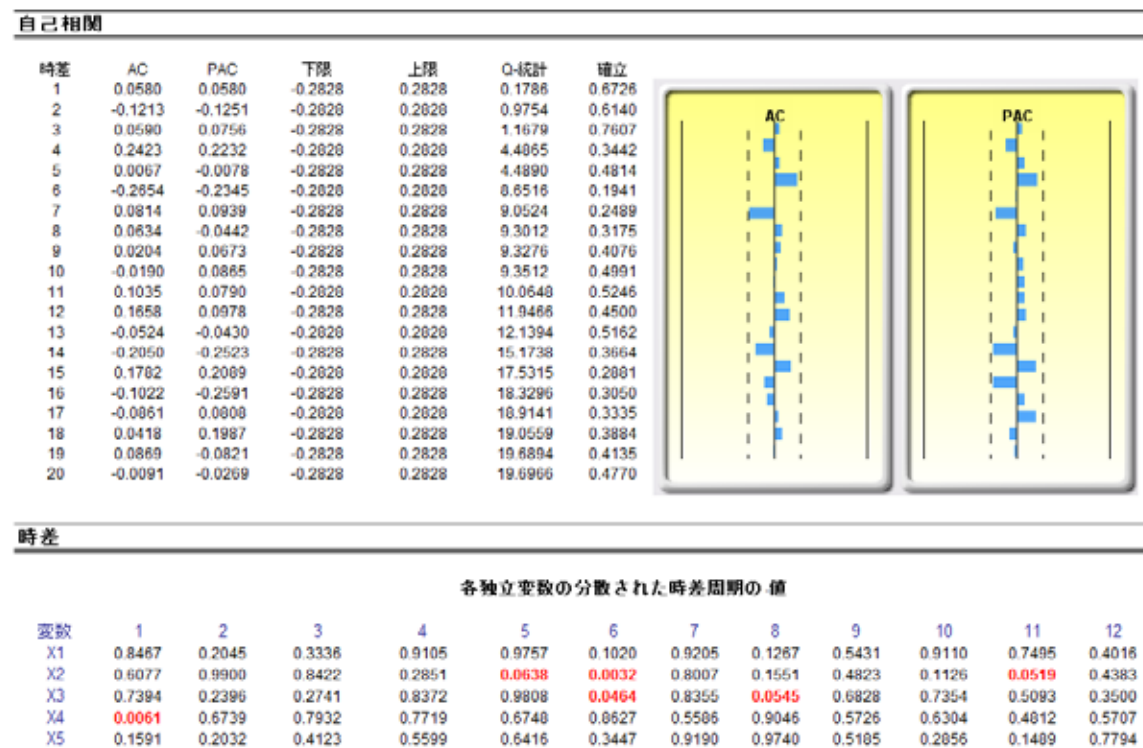


Figure 5.24 – 自己相関と分布的な時間差の結果

回帰モデルを実行するにあたっての他の要求は、正規性の仮定とエラーの球形です。正規性の仮定が違反、および異常値が存在する場合、線形回帰の適合の利点は、余り強力ではなく、また可能な情報検定ではなくなってしまいます。そして、これは線形の適合、

またはそうでないものの間の相違を意味します。エラーが独立していなく、正規分布されていない場合、データは自己相関されなければいけない、または非線形から、他のもっとも破壊的なエラーから影響されている事を示しています。エラーの独立性は不均一分散の検定で検出できます(Figure 5.25)。

実行されたエラーの正規性検定は、ノンパラメトリック検定で、サンプルが描かれた母集団の特定な形状についての仮定を作成せず、小さいサンプルデータの設定の分析を認めます。このテストは、正規分布された母集団から描かれたサンプルエラーがどれかと帰無検定を評価するのに対して、対立仮説のデータサンプルは正規分布されていません。もし、計算された D-統計が、様々な有意な値に対して D-クリティカル値と等しいか、上回っている場合、帰無仮説を放棄し、対立仮説（エラーは正規分布していません）を認めます。また、D-統計が D-クリティカル値よりも小さい場合は、帰無仮説を放棄（エラーは正規分布しています）しません。この検定は二つの蓄積的な周波数に及びます：一つはサンプルデータ設定から得られ、もう一つは、サンプルの平均値と標準偏差に基づいた論理的な分布から得られます。

正常性の検定とエラーの球形

回帰モデルを実行するための他の条件は正常性の仮定とエラーの球形です。もし、正常性の仮定が違反、または異常値がある場合、長所に適合した線形回帰テスト、または利用できる補助的なテストもっとも有力な方法とはいえないでしょう。そしてこれ自体が線形の適合が存在するかどうかを定める意味を持っています。もし、エラーが独立性でも正常分散でない場合、データが自己相関するか非線形から影響されているか、他に破壊的なエラーの存在を示していることとなります。エラーの独立性は、不均一分散検定で見出すことが出来ます（診断レポートを参照ください）。

エラーが発生上の正常性検定はノンパラメトリック検定で、描かれた母集団の詳細を気にしないため、仮定が無必要となり、小さなサンプルデータ設定から分析が実行できます。この検定は、データサンプルが正常に分散されていない対立仮説に対して、正常に分散された母集団から描かれたサンプルエラーがあるかどうかで帰無仮説を評価します。もし、いくつかの有意な値において計算されたD統計がDクリティカル値に同一、またはそれより上回っている場合、帰無仮説を棄却し対立仮説を認めてください（エラーは正常的に分散されていません）。一方、計算されたD統計がDクリティカル値より下回っている場合、帰無仮説を棄却しないでください（エラーは正常的に分散されています）。このテストは2つの累積度数に頼っています：一つはサンプルデータ設定から与えられており、二つ目は、平均値とサンプルデータの標準偏差に基づいた理論的な分散です。

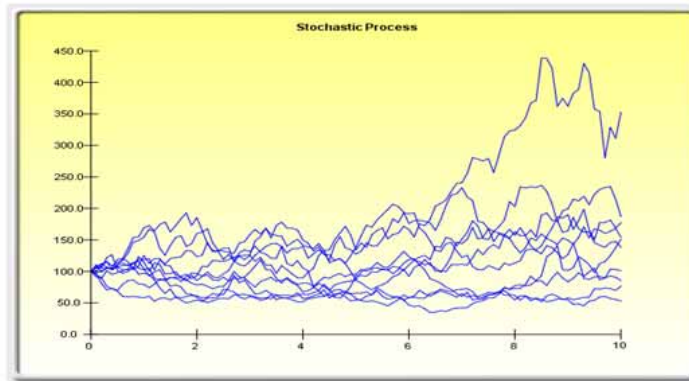
検定結果

		エラー	関係した周波数	観測	予期	O-E
平均回帰エラー	0.00	-219.04	0.02	0.02	0.0612	-0.0412
エラーの標準偏差	141.83	-202.53	0.02	0.04	0.0766	-0.0366
D 統計	0.1036	-186.04	0.02	0.06	0.0948	-0.0348
D クリティカルを 1%	0.1138	-174.17	0.02	0.08	0.1097	-0.0297
D クリティカルを 5%	0.1225	-162.13	0.02	0.10	0.1265	-0.0265
D クリティカルを 10%	0.1458	-161.62	0.02	0.12	0.1272	-0.0072
帰無仮説: エラーは正常分散されています。		-160.39	0.02	0.14	0.1291	0.0109
結果: エラーは正常分散されています。 1% のアルファレベルで。		-145.40	0.02	0.16	0.1526	0.0074
		-138.92	0.02	0.18	0.1637	0.0163
		-133.81	0.02	0.20	0.1727	0.0273
		-120.76	0.02	0.22	0.1973	0.0227
		-120.12	0.02	0.24	0.1985	0.0415

Figure 5.25 – エラーの正規性の為の検定

時々、あるタイプの時系列データは、確率過程以外の方法ではモデル化できません。これは、線で惹かれたイベントの性質は確率的だからです。例えば、モデルを適合できず、株式価格、利率、石油の価格と物価の価格をシンプルな回帰モデルを使って予測できないとします。これは、これらの変数は最も不確実で、揮発しており、前もって定められた振る舞いの静的なルールを辿らない為です。つまり、過程は季節性でないという事です。季節性は検定の実行の使用によって確認される一方、他の視覚系口は自己相関のレポートで検出されます(ACF は、ゆっくりとした下降の傾向を見せます)。確率過程は、イベントのシーケンス、および、確率法律によって生成された経路です。これは、ランダムイベントは、時間が経つに連れて発生できますが、特定の統計的、そして確率的なルールで支配されています。主な確率課程は、ランダムウォーク、またブラウン運動、平均回帰とジャンプ拡散が含まれています。これらの過程はランダム傾向を辿り、確率法則で限られた複数の変数を予測するのに使用されます。過程を生成する方式は良く知られていますが、現在の生成された結果は知られていません。(Figure 5.26)。

ランダムウォークブラウン運動の過程は、株式価格、物価と他の確率的な時系列データが与えるドリフト、または成長率とドリフト経路周辺の成長率を予測するのに使用できます。平均回帰の過程は、長期値を目的とした経路を認め、利率やインフレーションの比率（市場が取締権限による長期目的）のような長期の比率を持った時系列の変数の予測の為に使用やすくし、長期目的の比率のランダムウォーク仮定の分散を減少できます。ジャンプ拡 J 仮定は、石油の価格、または、電気の価格（分離した外因性のイベントの衝撃は価格を上昇するか下降します）の様な変数がランダムジャンプを時々表示できる時に時系列データの予測に適用します。これらの仮定は必要に応じて混合でき、適合する事が出来ます。



統計の概略

次に表示されているのは、提供されたデータがある確率過程のための推定変数です。これはユーザーの選択次第で、もし適合の確率が(適合の長所の計算に類似した)確率過程の予測の使用の保証が十分かどうかを定め、肯定の場合、そこにランダム・ウォーク、平均回帰、またはJumpDiffusionモデル、およびこれらの混合があるかどうかを定めてくれます。正しい確率過程モデルを選択すると、財政の予想で根本的なデータ設定が最もよく表される過去の経験と経済的なアプリオリに頼ることが出来ます。これらのパラメーターは確率過程の予測法に含むことが出来ます(シミュレーション (Simulation) | 予測 (Forecasting) | 確率過程 (Stochastic Processes))。

周期			
ドリフト率	-1.48%	回帰率	283.89%
ボラティリティ	88.84%	長期値	327.72
		飛翔率	20.41%
		ジャンプサイズ	237.89

確率も出る適合の確率: 46.48%
高い適合は、確立モデルの方が、習慣的も出るより良い事を示しています。

実行回数	20	基準標準	-1.7321
ポジティブ	25	P-値 (1-テール)	0.0416
ネガティブ	25	P-値 (2-テール)	0.0833
予期された実行	26		

低い p-値 (0.10, 0.05, 0.01より下)は、シーケンスはランダムではない為、季節性の問題を起こします。
また、ARIMAモデルの方が適切していることを示しています。一方、高い p-値は、ランダムで確率過程も出るが適切な事を示しています。

Figure 5.26 – 確率過程パラメーターの推定

多重共線性の存在は、線形的な関係が独立変数の間にある時に見られます。これが起きる時は、回帰の方式は完全に推定できません。共線性に近いシチュエーションでは、推定された回帰の方式はバイアスされ、不確実な結果を与えます。このシチュエーションは、ステップワイズの回帰が適用される時に特に発生し、統計的な有意な独立変数は、予想よりも早く回帰ミックスに投げ出され、回帰方式の結果として、どちらも効果的で確実です。重回帰方式の多重共線性の存在の一つの早い検定は、R-平方された値は比較的に高く、t-統計は比較的低いという事です。

もう一つの手早い検定は、独立変数の間の相関マトリックスを作成する事です。高い横断的な相関は、自己相関の強さを示しています。目分量は、相関と 0.75 より大きい絶対値で、厳しい多重共線性の表です。多重共線性の為の他の検定は、分散インフレーション要素 (VIF)の使用で、各独立変数を全ての他の独立変数に回帰して得る事で、R-平方

の値と VIF の計算を得る事が出来ます。2.0 を超える VIF は、厳しい多重共線性として考慮されます。10.0 を超える VIF は、破局的な多重共線性を示しています (Figure 5.27)。

相関マトリックス					
相関	X2	X3	X4	X5	
X1	0.333	0.959	0.242	0.237	
X2	1.000	0.349	0.319	0.120	
X3		1.000	0.196	0.227	
X4			1.000	0.290	

変動インフレーションの要因					
VIF	X2	X3	X4	X5	
X1	1.12	12.46	1.06	1.06	
X2	N/A	1.14	1.11	1.01	
X3		N/A	1.04	1.05	
X4			N/A	1.09	

Figure 5.27 – 多重共線性のエラー

相関マトリックスは、変数のペアの間でピアソンの商品モーメント相関 (一般的にピアソンの R として参照されている) として表示されています。相関係数は、-1.0 と +1.0 の間の範囲も含みます。符号は、変数の間の関係の方向を示し、係数は関係の強さと大きさを示します。ピアソンの R は、線形的な関係を測定し、非線形の関係の測定ではより少ない効果を示します。

どちらの相関が有意かを検定するには、両側-テールの仮説検定が実行され、結果としてなった p-値は上記のリストに表示されています。0.10, 0.05, と 0.01 より小さい p-値は、青で表示されており、統計的な有意性を示しています。つまり、相関のペアの為に与えられた有意な値より小さい p-値は、ゼロから統計的に異なっており、二つの変数の間に有意な線形な関係がある事を示しています。

二つの変数 (x と y) の間のピアソンの商品一次相関係数(R)は、共分散(cov)の測定に関連され、 $R_{x,y} = \frac{COV_{x,y}}{s_x s_y}$ と表示されます。共分散を二つの変数の標準偏差(s)で割る事の利点は、結果としなる相関係数は-1.0 と +1.0 の間に当たります。これは相関を良い比較的な測定とし、異なった変数の間の比較 (特に異なったユニットと大きさ)を与えてくれます。

スピアマンの非線形相関に基づいた順位も下記に含まれています。スピアマンの R は、ピアソンの R に関連されていて、データはまず順位別にランクされ、その後相関されます。ランクの相関は、二つの変数の一つ、または両方とも非線形の時、これらの関係はより良く推定されます。

有意な相関には原因は関連されないという事に重点を置いてください。変数の間の関係は、一つの変数は他の変数の変換に及ぼす原因となります。二つの変数がお互いに単独的に移動していますが、関連性の経路ではこれらは相関されますが、これらの相互関係は似せられなければいけません(例、サンスポットと株式マーケットの間の相関は強くなければいけませんが、一つは推量出来、原因が無く、これらの相互関係は純粹に似ています)。

他の強力なリスクシミュレーターのツールは統計分析ツールで、データの統計的な性質を定義します。診断の実行は、データの確率性質と検定の為に基本的な統計の記述による複数の統計的な性質のデータの確認が含まれています。

手順:

- 例証モデルを開き(リスクシミュレーター (**Risk Simulator**) | 例証 (**Examples**) | 統計的分析 (**Statistical Analysis**))、データワークシートと変数名称を含めたデータを選択してください(セル **C5:E55**)。
- リスクシミュレーター (**Risk Simulator**) | ツール (**Tools**) | 統計分析 (**Statistical Analysis**) (Figure 5.28) を選択してください。
- データタイプを確認し、選択されたデータは、列に変換された単一変数からかそれとも複数の変数からかのどちらかです。この例証では、選択されたデータの範囲は、複数の変数からだと考慮してください。終了後、**OK** をクリックして下さい。
- 希望する統計的検定を選択してください。お勧めは (デフォルトで) は、全ての検定を選択する事です。終了後、**OK** をクリックして下さい(Figure 5.29)。

実行された検定のより良い読解の為に生成されたレポートに目を通してください (Figures 5.30-5.33 でサンプルレポートが表示されています)。

Data Set (データー設定)

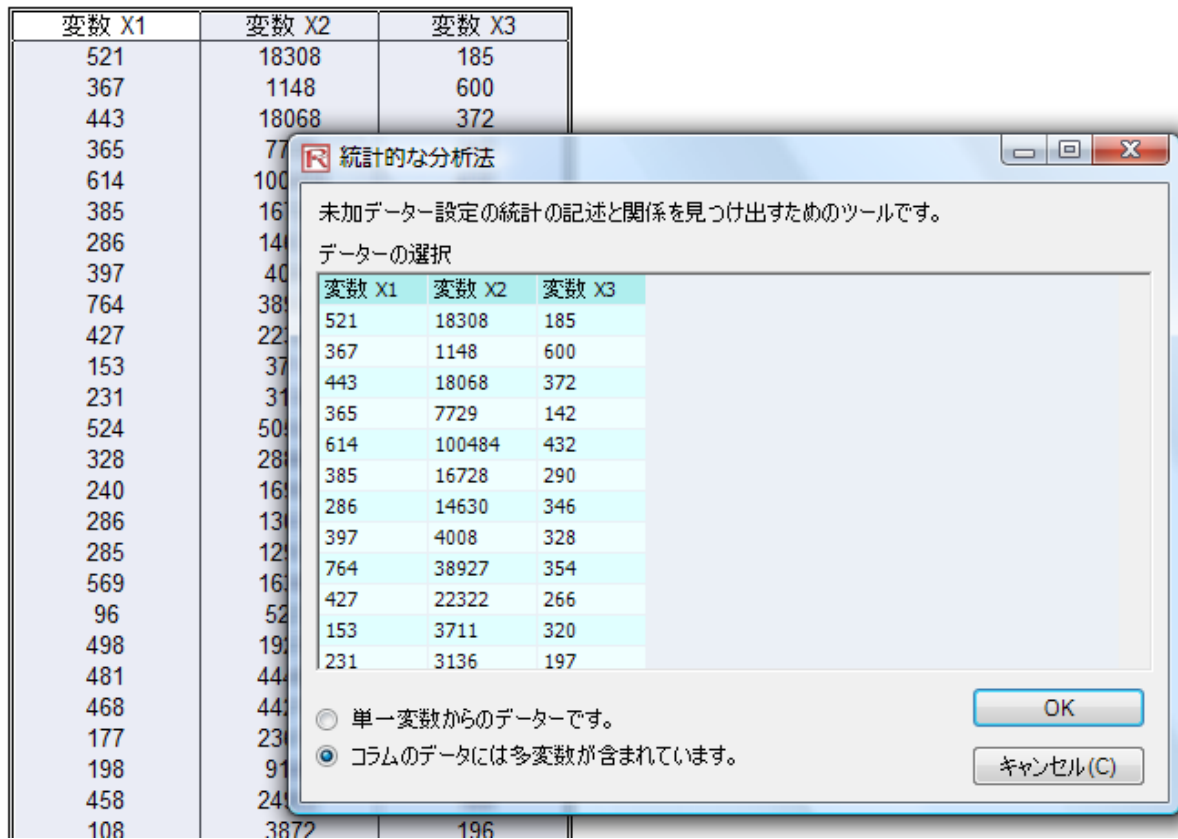


Figure 5.28 – 統計的分析ツールの実行

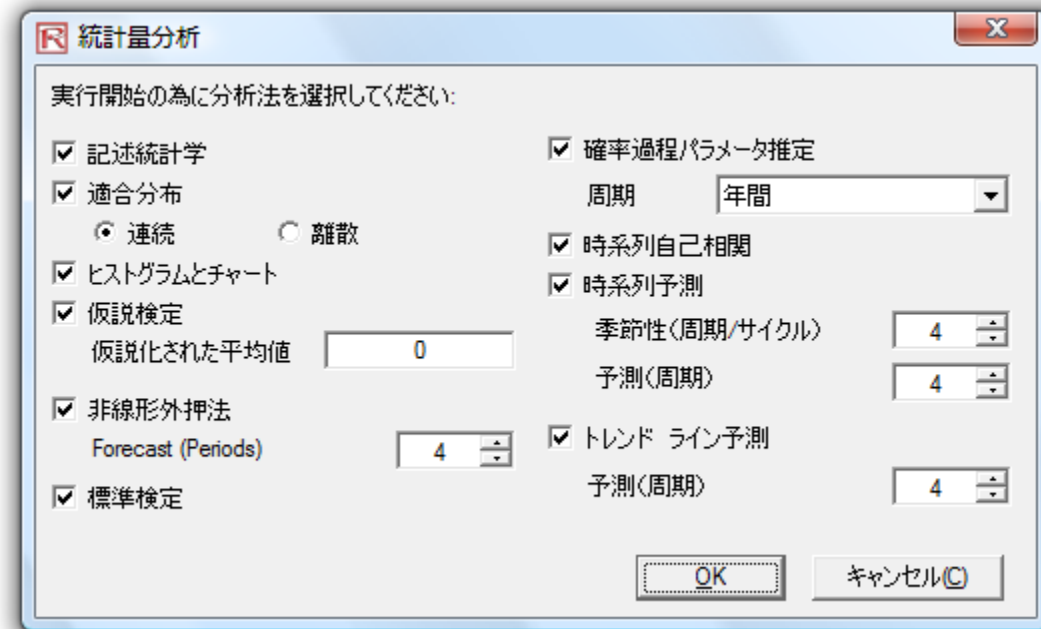


Figure 5.29 – 統計的検定

記述統計

統計の分析

ほとんど全ての分布法は4つのモーメントで記述できます（幾つかの分布法は一つのモーメントが必要とするのに対して、幾つかの分布法は二つのモーメントが必要となります。このようにして4つのモーメントを必要とする分布法もあります。）。記述的な統計量は、これらのモーメントを捕獲します。最初のモーメントは分配の配置を記述し（例、平均値、中間値とモード）、予期された値、予期されたリターン、または発生した平均値として表示されます。

算術平均値は、全てのデータポイントの足し算とこれらをポイント数で割ることで全ての発生した平均値を計算します。幾何学的な平均値は、全ての正のデータポイントの商品の力根を取ることで計算されます。幾何学的な平均値は、有意的に変動する比率やパーセントに対して最も確実です。例えば、変数の比率が含まれた複利から与えられた平均成長率を計算するのに使用できます。整えられた平均値は極異常値が整えられた後のデータ設定の幾何学的な平均値を計算します。平均値は、異常値が存在する時に有意的なバイアスの傾向を示し、整えられた平均値はこれらのバイアスを歪度の分配で減少します。

平均値の標準エラーは、サンプル平均値周辺のエラーを計算します。サンプルサイズが大きいく程、無限に大きいサンプルサイズのようにエラーが少なくなります。エラーがゼロと表示された場合は、母集団のパラメーターが推定されたことを示します。エラーをサンプル化するには95%の平均値として信頼区間が与えられます。サンプルのデータポイントの分析に基づき、現在の母集団の平均値は、平均値の区間の上下の間に当てはまります。

中間値は、全てのデータポイントの50%はこの値を上回り、残りの50%は下回るデータポイントです。最初の三つのモーメントの統計の間では中間値は異常値に対して少々見分けにくい傾向を見せます。対照的な分配は、中間値が算術平均値と同一しています。歪分配は、中間値が平均値からかなり離れているときに存在します。モードは、最も発生するデータポイントを測定します。

最小値はデータ設定内で最も小さい値を示し、最高値は、最も大きい値を示しています。範囲は最高値と最小値の間の違いを示しています。

二つ目のモーメントは分配の離散、および範囲を測定します。また、標準偏差、変動、4部位数、4部位数の中間の範囲などの測定で記述されます。標準偏差は、それぞれの平均値の全てのデータポイントの平均偏差を示します。リスクに関連されているように、とても一般的に測定です（高い標準偏差は、幅広い分布でリスクが高い事を意味し、または平均値周辺で広く離散されたデータポイントを示しています）。そして、これらのユニットは、オリジナルのデータ設定にとっても相対しています。サンプルの標準偏差は、小さいサンプルサイズに対して自由度の修正を計算に使用することから母集団の標準偏差と異なります。因みに、低、高信頼区間は標準偏差から与えられ、本当の母集団の標準偏差はこれらの区間無しで見出されます。データ設定が母集団の全ての要素をカバーしている場合、母集団の標準偏差を変わりに使用してください。二つの変動の測定は、標準偏差の平方された値です。

変動の係数は、サンプルの標準偏差をサンプルの平均値で割ったもので、離散のユニット・自由測定を証明し、様々な分布法の間で比較が出来ます（mとkg、何百万ドルと何億ドルの比較など）。最初の4部位数は、小さい値から大きい値に調整された時のデータポイントの25%を測定します。三つ目の4部位数は、データポイントの75%の値です。時々、異常値を無視する為に、データ設定を破損するように、分配の上下の範囲として4部位数は使用されます。4部位数の中間の範囲は、一つ目と三つ目の4部位数の間の違いを表示し、分配の中心の幅を測定するのに使用されます。

歪度は分配の三つ目のモーメントを表示します。歪度は、平均値周辺の分配の非対称の密度が性質です。正の歪度は、非対称的なテールの分布を示し、否の歪度は対照的なテールの分布を示しています。

尖度は、ある程度のピーク度、および正規分配に比べて分配の平坦を示しています。分配の4つ目のモーメントです。正の尖度値は、比較的にピークされた分配を示しています。否の尖度は、比較的な平坦な分配を示しています。ここで測定された尖度は、ゼロに整えられています（他の尖度の測定は3.0に整えられています）。両方が同一的に有効な間、ゼロに整えることは解釈を簡易にしてくれます。高い正の尖度は、中心周辺にピークのある分配を示し、急尖的または、太いテールを示します。これは、規制分布で予期されたよりも極的なイベントの高い確率を示しています（例えば、破局的なイベント、テロリストの攻撃、株式市場衝突など）。

統計の概略

統計	変数 X1		
観測	50.0000	標準偏差 (サンプル)	172.9140
算術平均	331.9200	標準偏差 (母集団)	171.1761
幾何学的な平均	281.3247	標準偏差への低信頼区間	148.6090
整えられた平均	325.1739	標準偏差への高信頼区間	207.7947
算術平均の標準エラー	24.4537	変動 (サンプル)	29899.2588
平均値への低信頼区間	283.0125	変動 (母集団)	29301.2736
平均値への高信頼区間	380.8275	変動の係数 Coefficient of Variability	0.5210
中間	307.0000	一周期 (Q1)	204.0000
モード	47.0000	3周期 (Q3)	441.0000
最小	764.0000	四分位数の中間の範囲	237.0000
最高	717.0000	歪度	0.4838
範囲		尖度	-0.0952

Figure 5.30 – 統計的分析ツールのレポートのサンプル

仮説過程 (一つの変数の平均母集団のt-検定)

統計の概略

データー設定からの統計		計算された統計	
観測	50	t-統計	13.5734
サンプル平均	331.92	P-値 (右側のテール)	0.0000
サンプルの標準偏差	172.91	P-値 (左側のテール)	1.0000
		P-値 (両側のテール)	0.0000
ユーザーが与えた統計		帰無仮説 (Ho): $\mu = \text{Hypothesized Mean}$	
仮説かされた平均値	0.00	対立仮説 (Ha): $\mu < > \text{Hypothesized Mean}$	
		注意: "<" 表示 "より大きい" 右側テール, "より小さい" 左側テール, および "同一ではない" 両側テールの仮説検定。	

仮説検定の概略

一つの変数の検定は、母集団の標準偏差値が知られていないが、サンプルの分配が正常に近似（統計はサンプルサイズが30より小さいが、最も大きいデーター設定と一緒に保守的な適切な結果を齎します）しているときに適切です。この検定は三つの仮説検定に適用できます：両側テールのテスト、右側テールテスト、左側テールテストです。三つの全てのテストとこれらの結果下記に表示されています。

両側テールの仮説検定

両側テールの仮説検定は、二つの変数の母集団が統計的に相似している帰無仮説 H_0 を検定します。また、対立仮説の場合は、サンプルデーター設定の使用で、仮説化された平均値から母集団が統計的に異なっている事を示します。計算されたp値が0.01、0.05、0.10より小、または同じの場合、仮説が拒絶された事を意味します。つまり、予測の平均値は統計的には、1%、5%と10%の有意水準に対して、かなり異なっているということです（または90%、95%と99%に対しての統計的な信頼度）。一方、p-値が0.10、0.05、および0.01より大きい場合は、母集団の平均値が、仮説化された平均値に統計的に相似していて、他の異なりのランダムチャンスが認められることを示しています。

右側テールの仮説検定

右側テール仮説は、母集団が統計的に仮説化された平均値に対して小さいか同一している帰無仮説 H_0 を検定します。また、対立仮説の場合は、サンプルデーター設定の使用で、仮説化された平均値に対して母集団が統計的に大きい事を示します。計算されたp値が0.01、0.05、0.10より小さいか、または同じの場合、仮説が拒絶された事を意味します。つまり、予測の平均値は統計的には、1%、5%と10%の有意水準に対して、かなり大きいということを示しています（または90%、95%と99%に対しての統計的な信頼度）。一方、p-値が0.10、0.05、および0.01より大きい場合は、母集団の平均値が、仮説化された平均値に統計的に相似しているか、小さいかを示しています。

左側テールの仮説検定

左側テール仮説は、母集団が統計的に仮説化された平均値に対して小さいか同一しているか、大きい帰無仮説 H_0 を検定します。また、対立仮説の場合

Figure 5.31 – 統計的分析ツールのレポートのサンプル(一つの変数の仮説検定)

正常性の検定

正常検定は、ノンパラメトリック検定の形状で、サンプルが描かれた母集団からの特定の形状に対して仮定を生成せず、小さいサンプルデーター設定の分析を認めます。このテストは、正常的に分配された母集団から描かれたサンプルデーターかどうかの帰無仮説を評価します。一方、対立仮説では、データーサンプルが正常的に分配されていません。もし計算されたp-値が、有意なアルファの値に同一か、より少ない場合、帰無仮説を放棄し対立仮説を認知します。したがって、p-値が有意なアルファレベルより高い間合い、帰無仮説を放棄しないでください。このテストは二つの蓄積した周波数に頼っています。一つ目は、サンプルデーター設定から、二つ目は、産婦津データーの標準偏差と平均値に基づいた理論的な分配からです。また、一つの方法として、正常性のための二乗平方検定があります。二乗平方検定は、ここで使用されている正常検定に比べて、実行するのにたくさんデーターポイントを必要とします。

検定結果

データー	相対的な周波数	観測	予期	O-E
平均データー	331.92			
標準偏差	172.91	47.00	0.02	0.02
D 統計	0.0859	68.00	0.02	0.04
1%のD基準	0.1150	87.00	0.02	0.06
5%のD基準	0.1237	96.00	0.02	0.08
10%のD基準	0.1473	102.00	0.02	0.10
帰無仮説: データーは正常に分配されています。		108.00	0.02	0.12
		114.00	0.02	0.14
結果: サンプルデーターは正常性 分散された		127.00	0.02	0.16
1%のアルファレベル、		153.00	0.02	0.18
		177.00	0.02	0.20
				0.1851
				0.0149

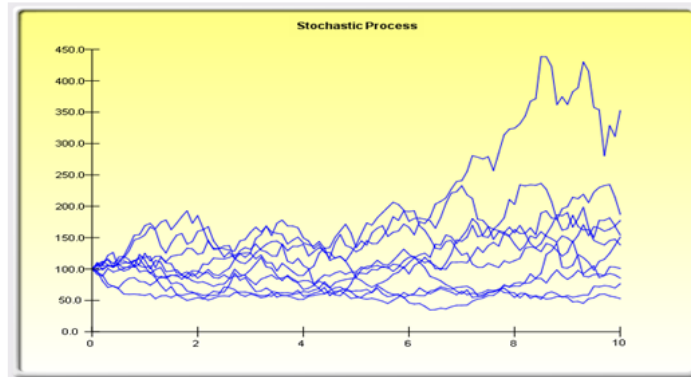
Figure 5.32 – 統計的な分析ツールのレポートのサンプル(正規性検定)

確率過程 - パラメーターの推定

統計の概略

確率過程は、確率法から生成された連続事象、または連続パッドの事を言います。これは、ランダムな事象は長期時間に発生する事がありますが、特定の統計的な、そして確率的なルールによって定められています。主な確率過程は、ランダムウォーク、またはブラウン運動、平均回帰とJumpDiffusionが、含まれています。これらの過程は、見た目ランダムな傾向を辿うが、確立法に制限されている複数の変数を予測する事が出来ます。このケースでの公式の生成過程は未知ですが、他のケースでは前もって知られています。

ランダムウォークブラウン運動過程は、株価、商品の価格、他の確立時系列データが齎したドリフト、または成長率とドリフトパッド周辺のボラティリティを予測する為に使用します。平均回帰過程は、長期の値を的にする為にパッドを許すことでランダムウォーク過程の変動を減少するために使用できます。これにより、利回り率やインフレーション率の長期率を持った時系列変数の予測を使いやすくします（これは、市場、または取締役権限 による長期の的の比率です）。JumpDiffusion過程は、変数が時折、オイルの価格や電気の価格などのランダム飛翔を表示したりする時系列データを予測するのに適切です（分離した外因性のイベントの衝撃は価格の飛翔の上昇または、下降することができます）。最後に、これらの三つの確率過程は混合、および必要に応じて調整することが出来ます。



統計の概略

次は、与えられたデータから得た確率過程の為に推定されたパラメーターです。適合の確立が（最良適合な計算に相似）確率過程の予測の使用の十分な保証がある場合、ランダムウォーク、平均回帰、およびjumpdiffusionモデル、またはこれらの混合が存在するかどうかを定義します。正しい確率過程のモデルの選択をすると過去の経験、以前のデータとどのデータ設定が代表的かの経済的な予想に頼ります。これらのパラメーターは、確率過程の予測に含まれます（シミュレーション（Risk Simulator） | 予測（Forecasting） | 確率過程（Stochastic Processes））。

ドリフト率	-1.48%	回帰率	283.89%	飛翔率	20.41%
ボラティリティ	88.84%	長期値	327.72	ジャンプ・サイズ	237.89
確率過程適合の確立性:		46.48%			

Figure 5.33 – 統計的分析ツールのサンプルレポート(確率パラメーターの推定)

分布的な分析ツール

これは、リスクシミュレーターでの統計的な確率ツールであり、多様な設定で有用であり、確率密度関数（PDF）の計算や離散分布（これらの名称を混合して使用します）では確率関数（PMF）にも使用でき、幾つかの分布とパラメーターが与えられている為、ある結果 x の発生確率を定める事が出来ます。また、累計分布関数(CDF)は計算でき、この x 値までの PDF の足し算です。最後に、反累計分布関数(ICDF)は、所与の発生での確率 x 値の計算に使用されます。

このツールは、リスクシミュレーター（Risk Simulator） | ツール（Tools） | 分布的な分析（Distributional Analysis）を選択する事で起動できます。例証のように、Figure 5.34

では、2 項分布の計算(例、コインを投げる等の 2 つの結果を持った分布、結果が表またはテールのどちらかで、これらの規定された確率がある)を表示しています。例えば、コインを 2 回投げるとし、結果の表を成功と設定し、2 項分布を試行回数 = 2 (コインを 2 回投げる)とし、確率を= 0.50 (表が出たら成功の確率)とします。PDF の選択と x 値の範囲の設定を 0 から 2 にし、ステップサイズを 1 (これは、x の値に 0, 1, 2 を要求している事を意味します)とし、分布の論理的な 4 つのモーメントのように結果として成った確率は表とグラフに与えられます。投げたコインの結果として表-表とした場合、裏-裏、表-裏と裏-表で、表を出さない確率は 25%となり、表一回は、50%と表 2 回は 25%です。

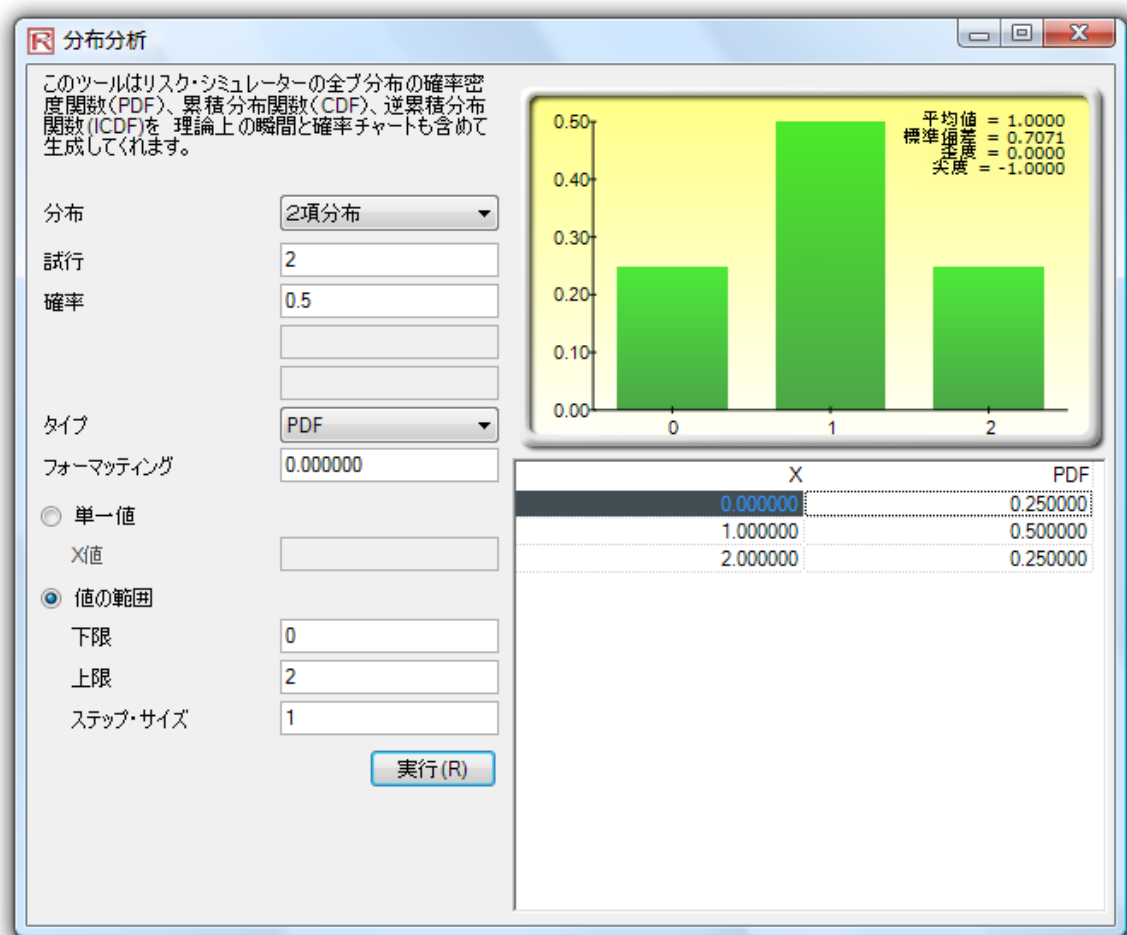


Figure 5.34 – 分布的な分析ツール(2 項分布と 2 の試行回数)

同様に、コインを投げる厳密な確率を得る事が出来ます。Figure 5.35 で表示されているように、20 の試行回数を行います。結果は、表とグラフの両方に記述されています。

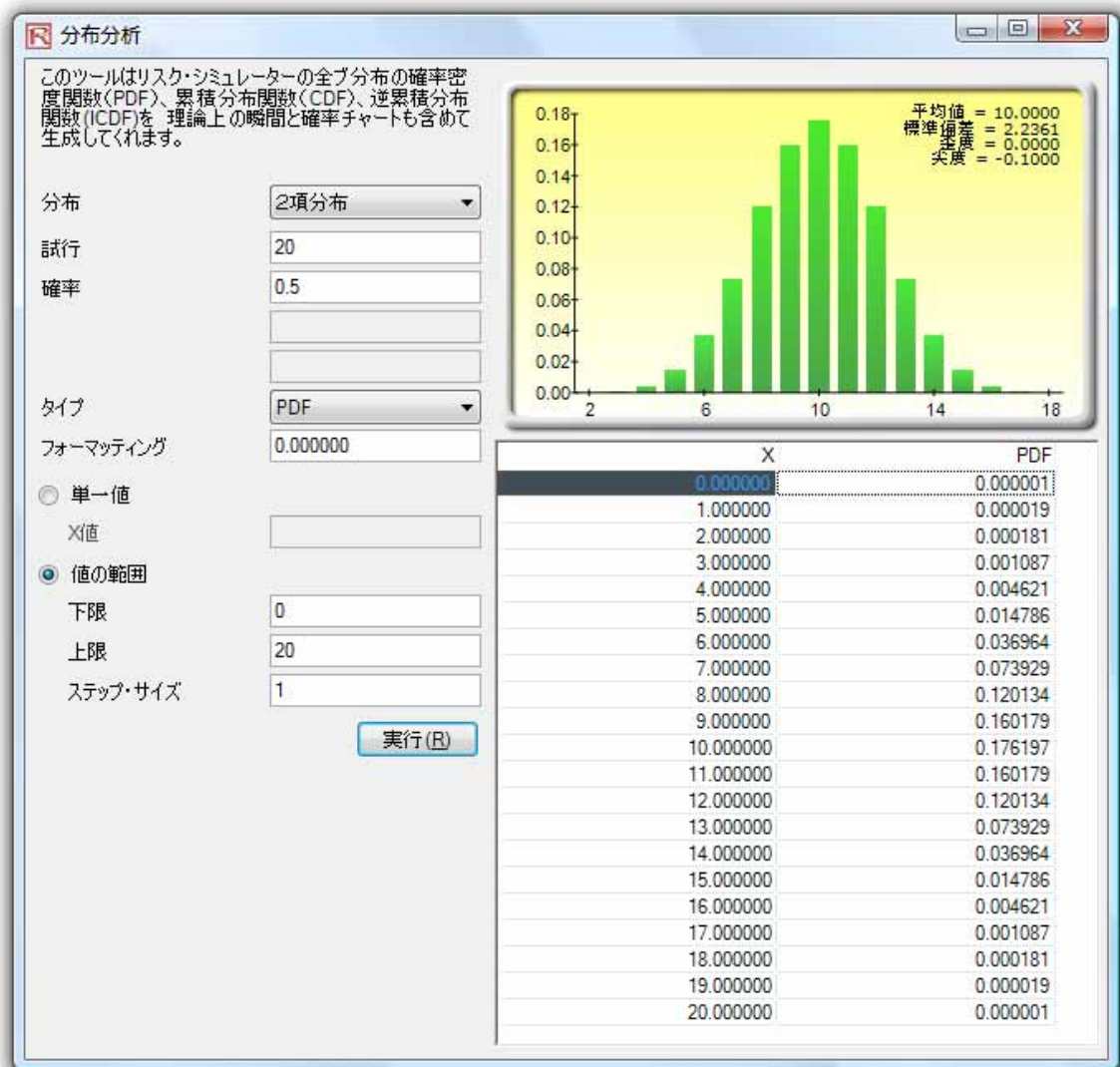


Figure 5.35 –分布的な分析ツール(2項分布と20の試行回数)

Figure 5.36 は、同じ2項分布を表示していますが、ここではCDFが計算されています。CDFとは、単にポイントxまでのPDFの値の足し算です。例えば、Figure 5.35では、0.1、と2の確率は0.000001, 0.000019, と0.000181である事が分かり、これらの足し算は0.000201で、x=2のCDFの値はFigure 5.36で表示されています。PDFが2つの表を得る事の確率(およそ0.1、と2の表の確率)を計算します。補分布は、(例、 $1 - 0.00021$ は0.999799、およそ99.9799%を得る)は、3つの表、およそそれ以上を得る確率を与えてくれます。

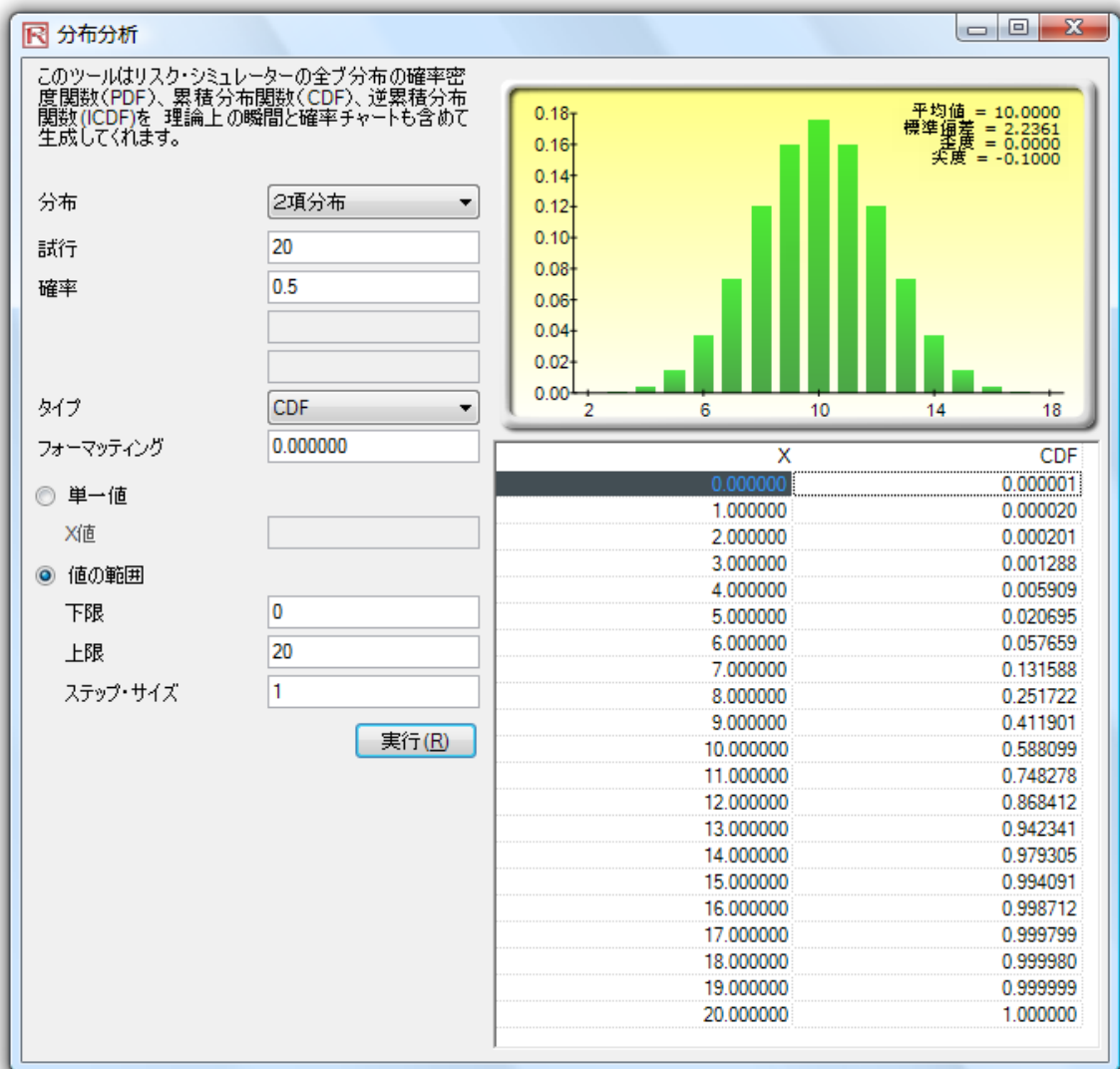


Figure 5.36 – 分布的な分析ツール(2 項式分布の CDF と 20 の 試行回数)

この分布的な分析ツールの使用によって、リスクシミュレーターでガンマ、ベータ、負の 2 項分布のような様々の高度な分布の分析ができます。連続的な分布と ICDF 機能でのツールの使用の例証以上に、Figure 5.37 は標準正規分布（平均値がゼロで標準偏差 1 の正規分布）を表示し、ICDF の適用によって、97.50%（CDF)の累積確率に当たる x の値を見出します。これは、97.50%の片側 CDF が、両側信頼区間の 95%（右側の確率として 2.50%、左側の確率として 2.50%に相当し、95%を中心および信頼区間の範囲に残し、一つのテールの範囲に 97.50%が同等します）に等しいという事です。結果は、よく知られた 1.96 の Z-スコアです。したがって、分布的な分析ツール、他の分布の為に標準スコア、他の分布の累計確率、および一致の使用は、早くそして簡単に得る事が出来ます。

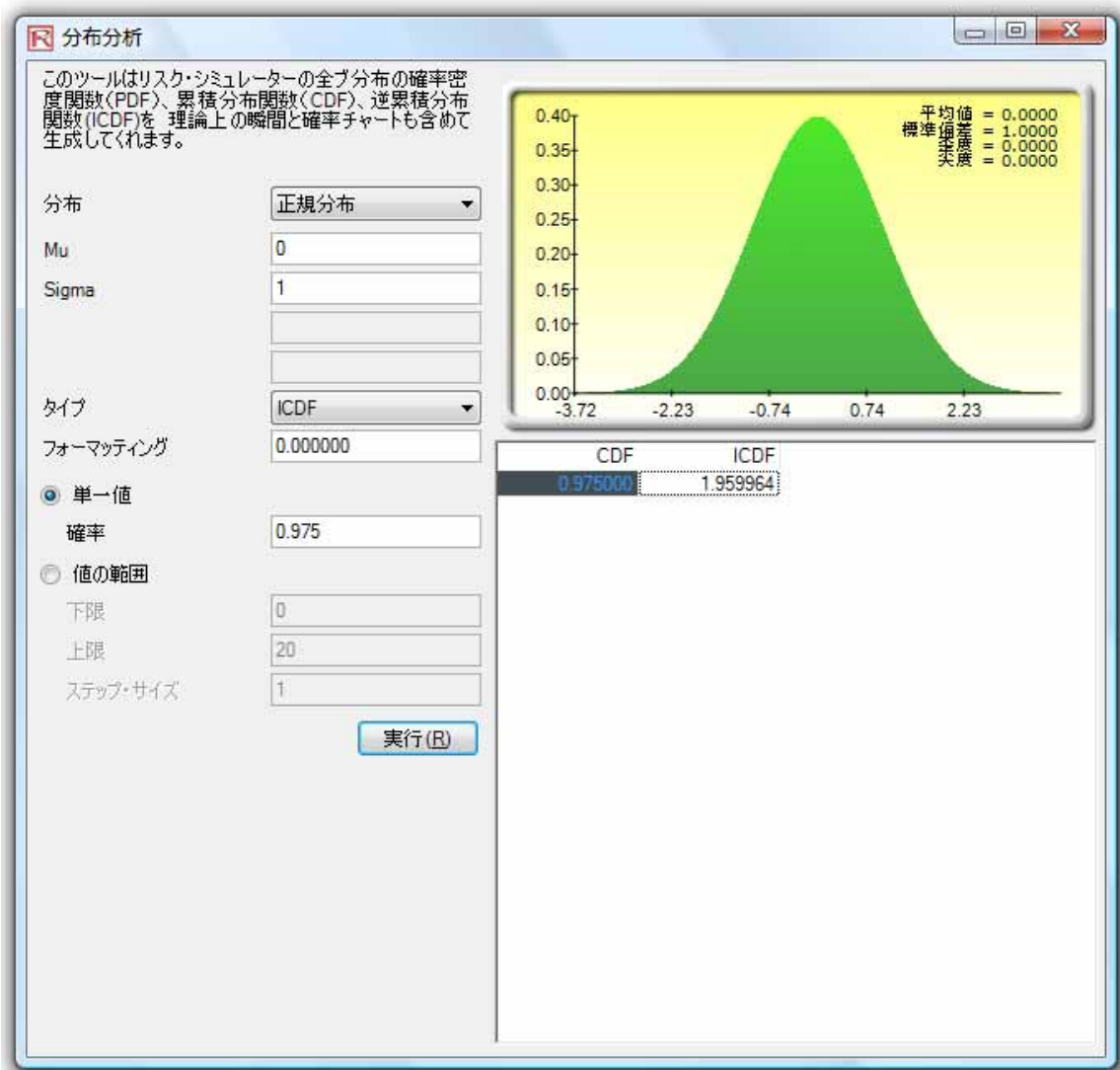


Figure 5.37 – 分布的な分析ツール (正規分布の ICDF と Z-スコア)

リスクシミュレーター2011/2012の新しいツール

乱数の生成、モンテカルロ 対 ラテン・ハイパーキューブとコピュラ相関法
2011/2012 版をはじめるにあたって、6 つの乱数の生成、3 つのコピュラ相関と 2 つのシミュレーション・サンプリング法の選択が可能となりました (Figure 5.41)。これらの選好は、**リスクシミュレーター/ オプション**を通して設定する事ができます。

乱数の生成 (RNG) は、どのシミュレーション・ソフトウェアでも中心にあります。様々な数学的分布が乱数の生成に基づくことで構成することができます。デフォルトな方法は、ROV リスクシミュレーター専売特許の技法であり、もっとも最適で強硬な乱数を提供します。サポートされた 6 つの乱数の生成があり、一般的に ROV のリスクシミュレーターのデフォルト法と高度な減算乱数シャッフル法の 2 つは、お勧めの使用アプリケーションです。他の技法は、モデル、あるいは分析が特別な使用を求めない限り適用しないでください。我々は、反対にこれらのお勧めの 2 つのアプローチで結果を検定することを進めています。2 つのアプローチとは、RNG の最下位リストの実行の速いシンプルなアルゴリズムと結果がより広範な RNG の最上位リストです。

相関セクションでは、3 つの技法がサポートされています: 正規コピュラ。T コピュラと擬似正規コピュラ。これらの技法は、数学的な統合テクニックに従属し、疑問がある場合、正規コピュラは、最も確実に保守的な結果を提供します。T コピュラは、シミュレーションされた分布のテールの極値へ提供し、擬似正規コピュラはこれらの値の間に相当する結果を提供します。

シミュレーション法のセクションでは、モンテカルロ・シミュレーション (MCS) とラテン・ハイパーキューブ・サンプリング (LHS) 法がバックアップされています。コピュラと他の多重関数は LHS と互換性を持っていないことに注目してください。これは、LHS は単一ランダム変数には適用できるが、共同の分布には適用できないということです。実際には、LHS は、モデル結果の確実性上にはとても有限されたインパクトを見せ、LHS からモデルにより分布があるほど、個々の分布にしか適用できないということです。最初に指名したサンプル数を完了しなければ LHS の用益は失われてしまいます。例えば、シミュレーション中のシミュレーション実行の停止です。また、LHS は、分布からの第 1 次サンプルの実行を行う以前に各分布からのサンプルを生成し、整理する必要がある多数の入力を含んだシミュレーションモデルの負担も適用することができます。これは大きなモデルの実行の著しい遅れの原因となる可能性があるが、やはり本の少しの確実性しか提供してくれません。最後に、分布が良い振る舞いを持っており、対照的で相関を持っていない場合に、最的な適用が行えます。それでも尚、LHS は、強力なアプローチであり、一様的にサンプル化された分布を提供します。LHS を適用すると他のより一様的にサンプル化された分布 (分布の全ての部分をサンプル化することができます) に

比べて、MCS は時々鈍重な分布を生成することがあります(サンプル化されたデータは、時々分布の1つのエリアに主に集中することがあるため)。

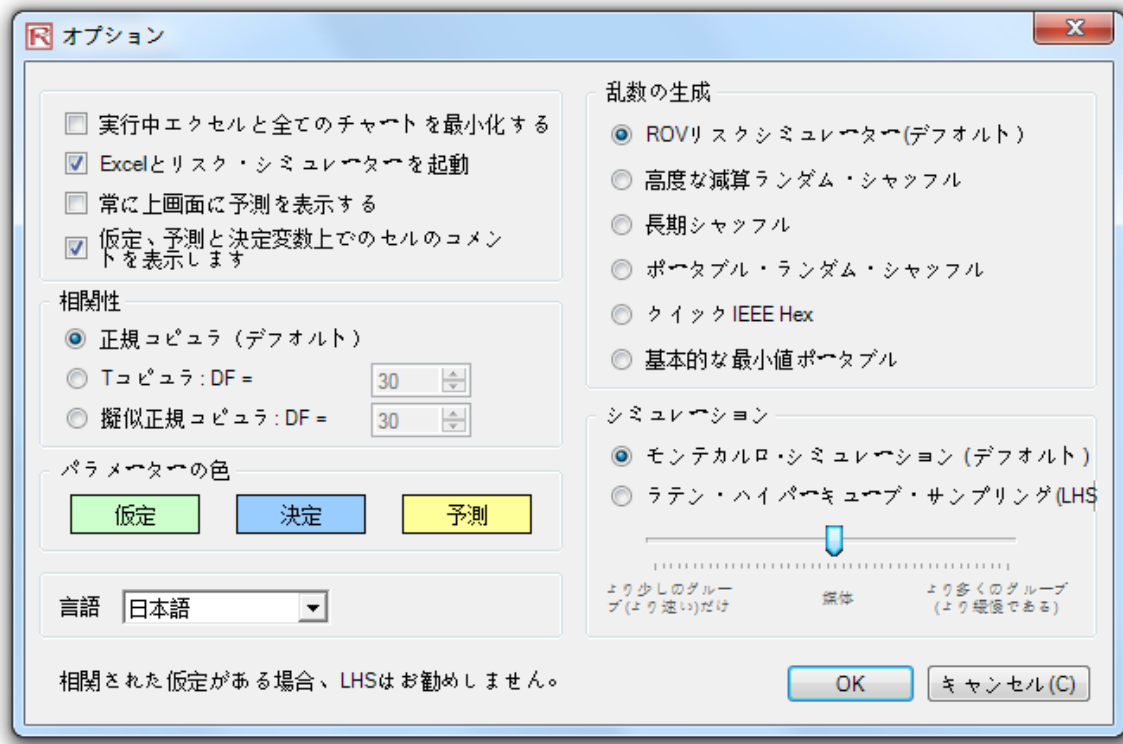


Figure 5.41 - リスクシミュレーターのオプション

非季節性データと傾向除去データ

このツールは、元のデータを無季節化と傾向除去をすることで、すべての季節性と傾向要素を取り出します(Figure 5.42)。予測モデルの過程では通常、季節性と傾向からの蓄積されたデータセットの効果の除去が含まれており、値の絶対的な変化だけを表示し、時系列データの設定の季節性サイクルの一般的なドリフト、傾向、ツイスト、バンドと影響を除去した後に強力なサイクルパターンの識別を可能にします。例えば、傾向除去のデータセットは、明確に指定された年の企業の売り上げのより確実な計算が必要になるかもしれません。これは、データセット全体のスロープをフラット面にシフト化することによって、強調されたサイクルと振る舞いのより良い観測ができます。

多くの時系列データは季節性を示し、特定のイベントはいくつかの期間周期、あるいは季節性周期後に反復を始めます(例、スキーリゾートの収益は、夏よりも冬の方が高く、この当たり前のサイクルは毎年冬に反復されることになります)。季節性周期は、反復以前に経過しなければならない周期数を示しています(例、一日に24時間、一年に12ヶ月、一年に4半期、一時間に60分など他)。このツールは、元のデータを非季節化、傾向除去し、すべての季節的な要素を取り出します。1を上回る季節性の目録は、季節性サイクルの中の高い周期、あるいはピークを示し、1を下回る値は、サイクルの沈下を示します。

過程 (非季節性と傾向除去):

- 分析するデータの選択 (例、B9:B28)。そして **リスクシミュレーター / ツール / データの非季節性と傾向除去** をクリック。
- **非季節性データ** と/あるいは **傾向除去データ** を選択し、実行したい傾向除去モデルのどれかを選択し、重要な順番に入力していき (例、多項式順、移動平均順、相違順と比率順)、**OK** をクリック。
- 技法、アプリケーション、結果グラフと非季節化／傾向除去化されたデータ上で生成されたより詳しい2つのレポートを確認。

過程 (非季節化と傾向除去化):

- 分析するデータの選択 (例、B9:B28)。そして **リスクシミュレーター / ツール / データの季節性検定** をクリック。
- 検定には、最大の季節性周期を入力。これは、6 と入力した場合、ツールは、次の季節性周期を検定することを示しています: 1, 2, 3, 4, 5, 6. 周期 1 は、もちろん、データに季節性がないことを示しています。
- 技法、アプリケーション、結果グラフと非季節化／傾向除去化されたデータ上で生成されたより詳しい2つのレポートを確認。最良の季節性周期が最初に一覧され (RMSE のエラー測定が低い順)、全ての重要なエラー測定は、比較のために含まれています: 2 乗平均平方根誤差 (RMSE)、平均二乗誤差 (MSE)、平均絶対偏差 (MAD) と平均絶対誤差率 (MAPE)。

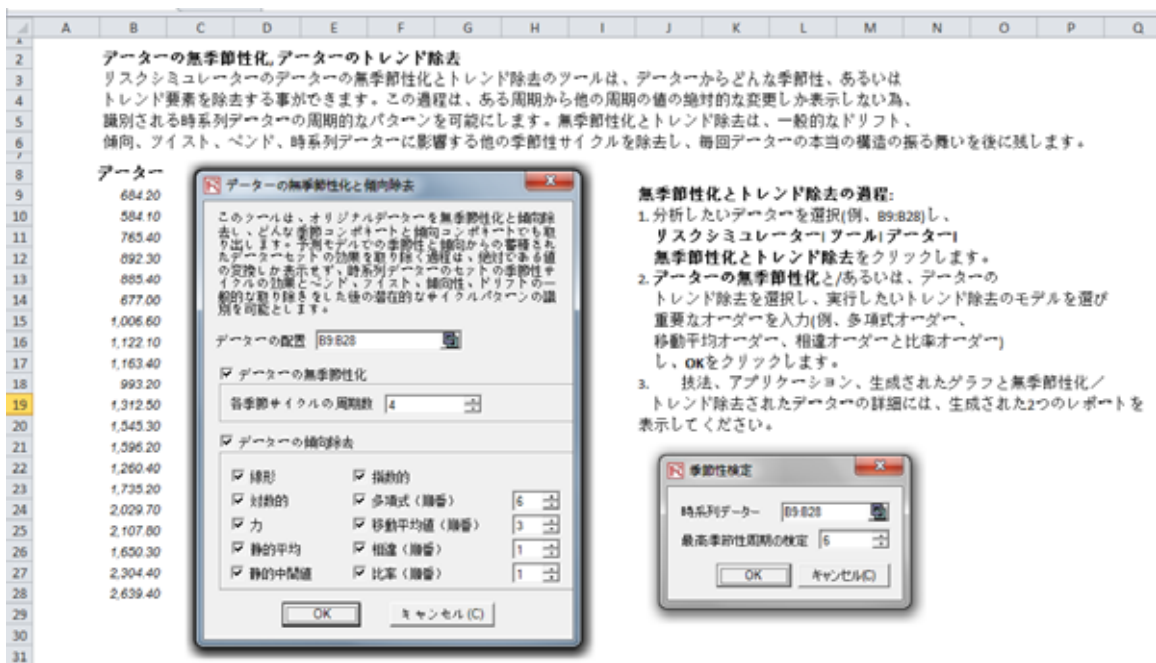


Figure 5.42 - 非季節性と傾向除去のデータ

主成分分析

主成分分析は、データのパターンを識別する方法で、相似や相違を強調するためのデータの作り直しをします(Figure 5.43)。データのパターンは、多次元で検出するのはとても困難で、複数の変数と多次元的なグラフがある場合の中断と象徴はとても難しくなってしまいます。データのパターンが一度検出されると、圧縮が可能となり、次元数を減らすことができます。このデータ次元の減少は、情報の損失に相当しません。代わりに、少ない変数の数でも同じ品質の情報が得られることを示しています。

過程:

- 分析するデータを選択(例、B11:K30)。 **リスクシミュレーター/ ツール / 主成分分析** を選択し、OK をクリック。
- 計算された結果で生成されたレポートの確認。

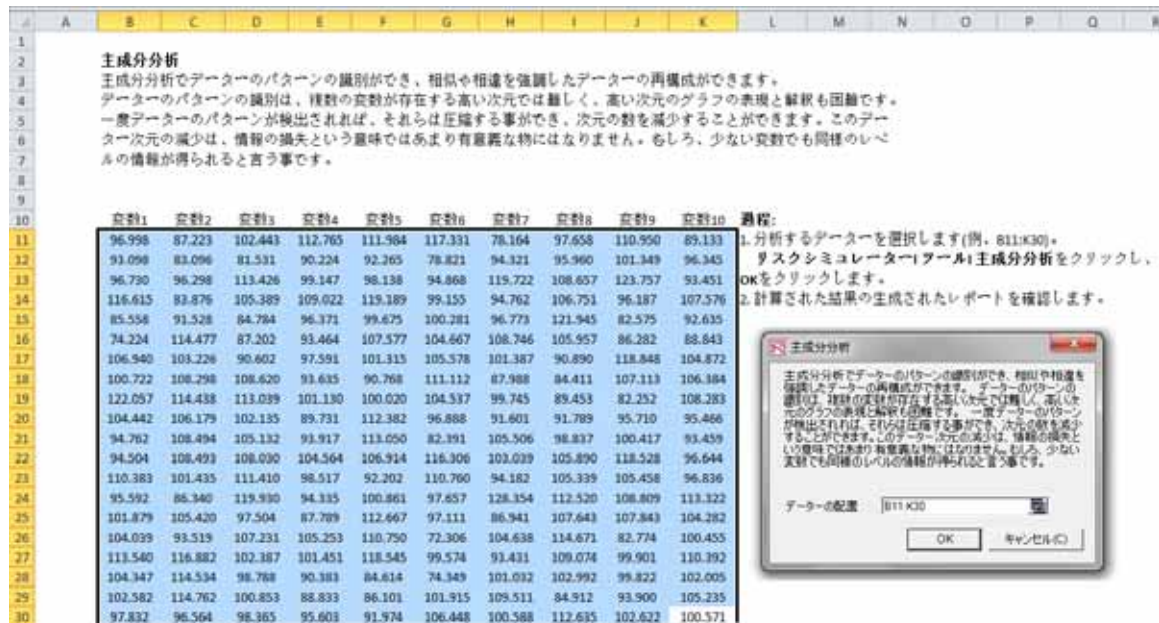


Figure 5.43 - 主成分分析

構造変化分析

構造変化検定は、異なったデータセットの係数が同じで、一般的に構造変化がある場合に時系列分析に使用されています(Figure 5.44)。時系列のデータセットは、2つのサブセットに分けることができ、各サブセットは、相互的に検定され、統計的に全てのデータセットに、特定の期間周期に変化の開始が本当にあるかどうかを定義します。構造変化検定は、新しいマーケティング・キャンペーン、事業、大イベント、獲得、権利奪取の検定のような母集団の異なったサブグループ上に独立変数が異なった影響を表示するかどうかを定義するためにたいいてい使用されます。データセットが 100 の時系列データポイントを持っていると仮定した場合、いくつかのブレイクポイントを設定する事ができます。例えば、10、30 と 51 のデータポイント(これは、次のデータ

ポイント上で3つの構造変化検定が実行されることを示しています：10-100 と比較されたデータポイント 1-9；30-100 と比較されたデータポイント 1-29；51-100 と比較されたデータポイント 1-50。本当に 10、30 と 51 のデータポイントの始めに強行された構造に変化があるかどうかを検定します)。片側テールの仮説検定は、2 つのデータのサブセットが統計的に、相互的に相似しているかどうかを帰無仮説 (H_0) 上で実行されます。これは、統計的に有意な構造変化がないことを示します。対立仮説 (H_a) は、2 つのデータのサブセット統計的に、相互的に相違しているかどうかを実行し、構造変化の可能性を示します。計算された p 値が 0.01、0.05、あるいは 0.10 と同じかそれ以下の場合、仮説が拒絶されたことを意味し、2 つのデータのサブセットが統計的に、そして有為的に 1%、5%と 10%の有意度で相違している事を示します。高い p 値は、統計的に有意な構造変化がないことを表しています。

過程:

- 分析するデーターを選択 (例、B15:D34)。**リスクシミュレーター/ ツール / 構造変化検定** をクリックし、データー上の適用したい重要な検定ポイントを入力し(例、6、10、12)、OK をクリック。
- データー上で統計的に有意なブレイクポイントを示す検定ポイントがどれかを定義するためにレポートを確認。

Y	X1	X2
521	18308	185
367	1148	600
443	18068	372
365	7729	142
614	100484	432
385	16728	290
286	14630	346
397	4008	328
764	38927	354
427	22322	266
153	3711	320
231	3136	197
524	50508	266
328	28886	173
240	16996	190
286	13035	239
285	12973	190
569	16309	241
96	5227	189
498	19235	358

過程:

1. 分析するデーターを選択(例、B15:D34)し、**リスクシミュレーター| ツール| 構造ブレイク検定** をクリックし、データー上の適用したい主な検定ポイントを選択 (例、6, 10, 12) し、OK をクリックして下さい。
2. これらの検定ポイントがデーター上の統計的に有意なブレイクポイントを示しているか、それとも示していないかを定義する為にレポートを表示してください。

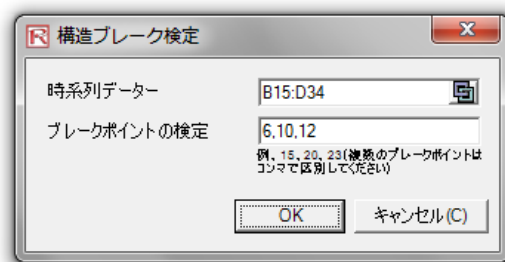


Figure 5.44 - 構造変化分析

傾向ラインの予測法

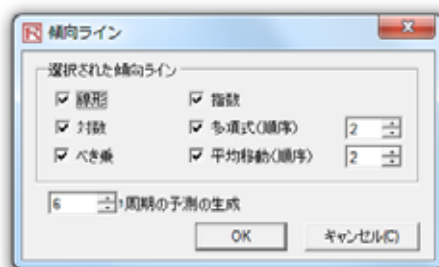
傾向ラインは、時系列データーが目立った傾向 (Figure 5.45) を辿っているかどうかを定義するために使用されます。傾向は、線形、あるいは非線形のどちらかに分類できます (指数、対数、移動平均、ベキ、多項式の様に)。

過程:

- 分析するデーターの選択し、**リスクシミュレーター/ 予測法/ 傾向ライン**をクリックし、データー上で適用したい重要な傾向ラインを選択します(例、デフォルトで全ての技法を選択することが可能)。予測する周期数を入力(例、6 周期)し、**OK**をクリックします。
- これらの傾向ラインの検定からデーターに最適、そして最も適合な予測を提供するのはどれかを定義するためにレポートを確認。

履歴的な販売収入

年	4 部位数	周期	販売
2006	1	1	\$584.20
2006	2	2	\$584.10
2006	3	3	\$765.40
2006	4	4	\$892.30
2007	1	5	\$885.40
2007	2	6	\$877.00
2007	3	7	\$1,006.60
2007	4	8	\$1,122.10
2008	1	9	\$1,163.40
2008	2	10	\$992.20
2008	3	11	\$1,312.50
2008	4	12	\$1,545.30
2009	1	13	\$1,596.20
2009	2	14	\$1,260.40
2009	3	15	\$1,735.20
2009	4	16	\$2,029.70
2010	1	17	\$2,107.60
2010	2	18	\$1,650.30
2010	3	19	\$2,304.40
2010	4	20	\$2,639.40



手順

このモデルを実行するには次の手順を行ってください。

1. 履歴的なデーターを選択してください (セル E25:E44)。
2. シミュレーション (Risk Simulator) | 予測 (Forecasting) | 傾向ライン (Trendlines)。

Figure 5.45 - 傾向ラインの予測法

モデルの確認ツール

モデルの作成後と過程と予測の設定後にいつものようにシミュレーション、あるいはモデルの確認ツール (Figure 5.46) を実行することができ、モデルが正確に設定されたかどうかを検定することができます。反対に、モデルが実行せず、設定のどこかが間違っていると思う場合は、このツールを [リスクシミュレーター / ツール / モデルの確認](#) から実行してください。モデルの何処を修正しなければならないのかが確認できます。このツールは、最も一般的なモデルの問題点やリスクシミュレーターの仮定と予測法を指摘を行うため、全ての問題解決に役立つとは限りません。モデルが正確に作動するように進展を維持する事が我々の課題でもあります。

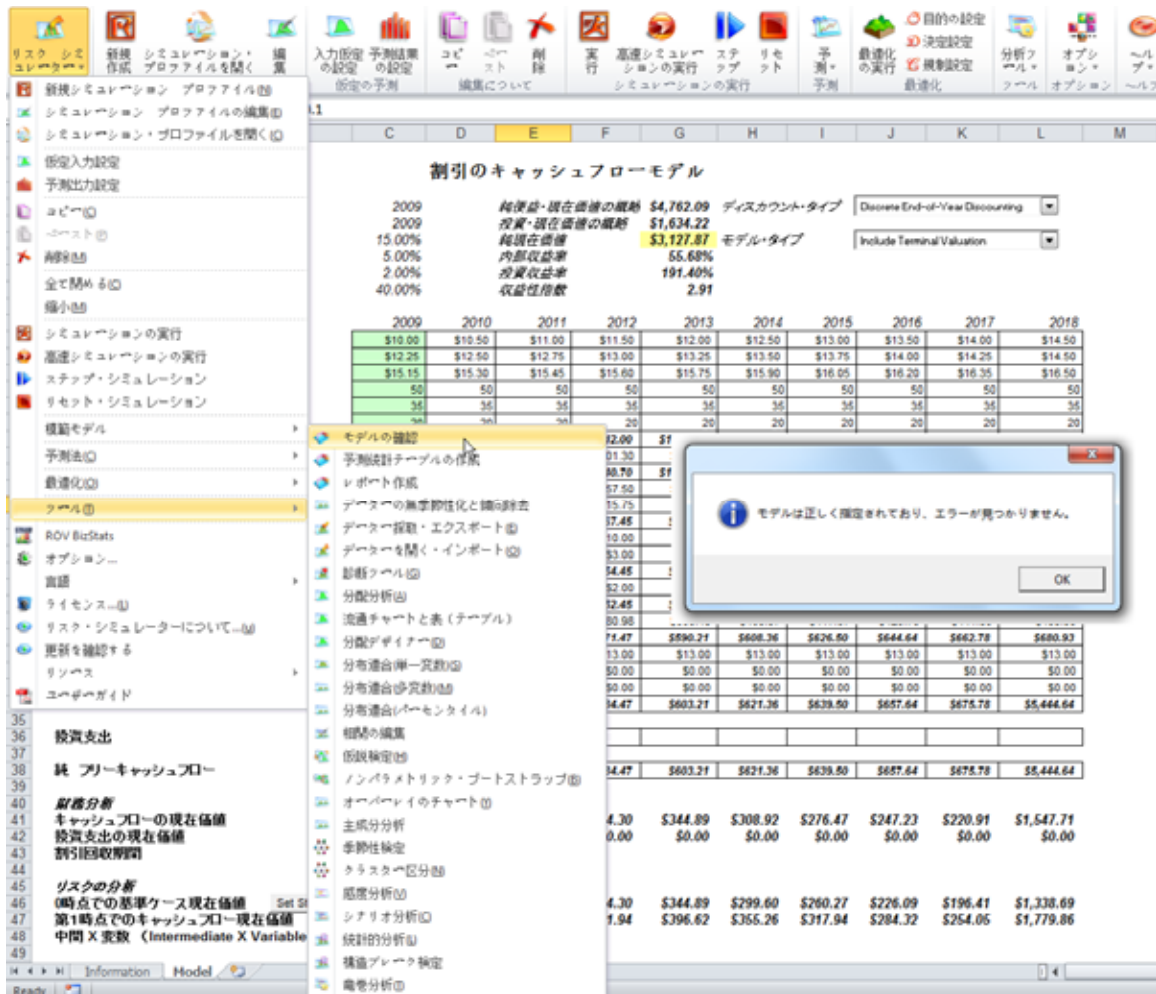


Figure 5.46 - モデルの確認ツール

分配的なパーセンタイル適合ツール

分配的なパーセンタイル適合ツール(Figure 5.47)は、も 1 つの確率分布の適合法です。関連した様々なツールがあり、それぞれの使用法と利点があります。

- 分配的な適合(パーセンタイル)--始めから対立技法を使用(パーセンタイルと第 1 次/2 次モーメントの組み合わせ)することで、未加データの必要なしで特定の分布の適合名パラメーターを検出することができます。この技法は、パーセンタイルとモーメントだけが利用可能で十分なデータがない時、あるいは、2 つや 3 つのデータポイントだけで定義された、あるいは知られている分布タイプで分布全体の回復を行う際に便利です。
- 分配的な適合(単一変数)--統計的な技法の使用を通して、最的な分布と入力パラメーターを見出すために 42 全ての分布に未加データを適合します。最的な適合には複数のデータポイントは必要ですが、分布タイプは事前に知る必要がありません。
- 分配的な適合(複数の変数)--統計的な技法の使用を通して、同時に複数の変数上に未加データを適合します。単一変数の適合と同様のアルゴリズムを使用しますが、変数の間にペアワイズ相関マトリックスを統合します。最的な適合には、複数のデータポイントが必要となりますが、分布タイプは事前に知る必要がありません。
- カスタム分布(仮定の設定)--ノンパラメトリック・リサンプリング法の使用を通して依存するデータでカスタム分布を生成し、実践的な分布に基づいた分布をシミュレーションします。少ないデータポイントが必要とされますが、分布タイプは、事前に分かりません。

過程:

- **リスクシミュレーター/ ツール / 分配的な適合(パーセンタイル)**をクリックし、使用したい確率分布と入力タイプを選択し、パラメーターを入力して**実行**をクリックして結果を得ます。適合された R2 乗結果を確認し、経験論 対 理論的な適合の結果を比較して、自分の分布が最適なのかを定義します。

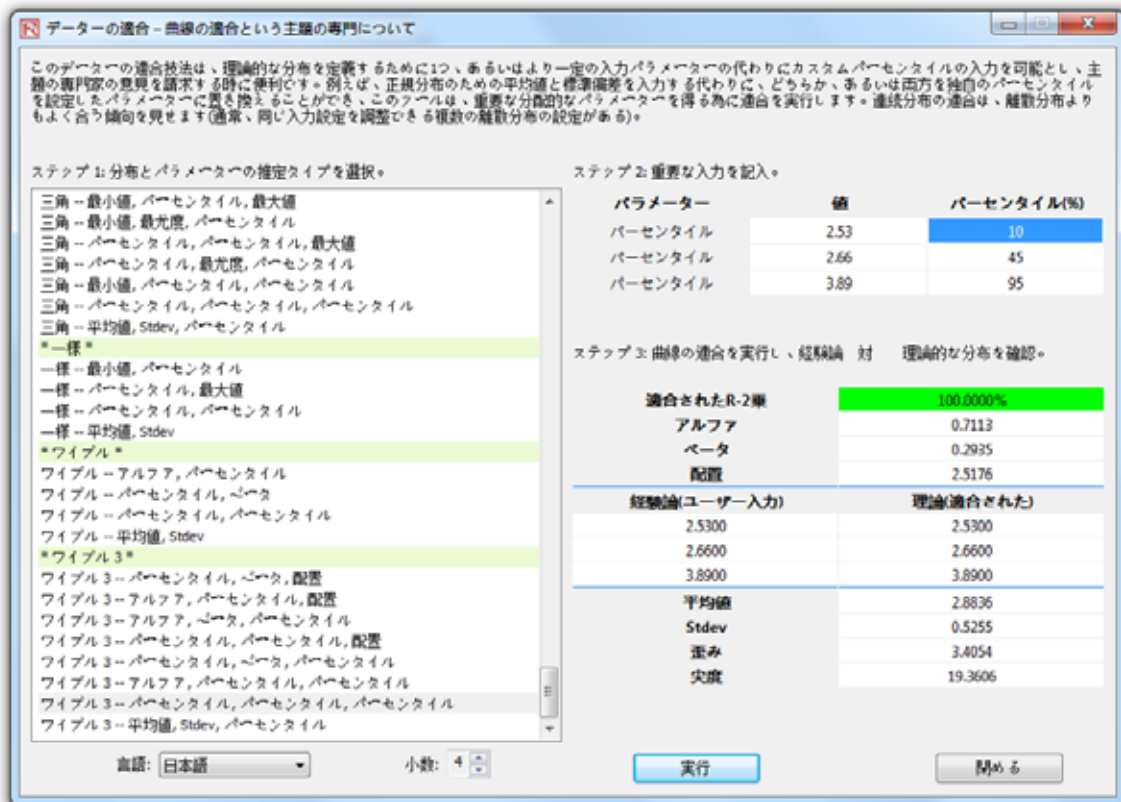


Figure 5.47 - パーセンタイルの分配的な適合ツール

分布グラフと表: 確率分布ツール

この新しい各位率分布ツールは、とても強力で速いモジュールであり、分布グラフと表の生成に使用されます(Figures 5.48-5.51)。リスクシミュレーターには、3つの似たようなツールがありますが、それぞれがとても異なった適用性を持っていることに注目してください。

- 分配的な分析--リスクシミュレーターで使用可能な42の確率分布のPDF、CDFとICDFの速い計算を行い、これらの値の確立表を提供します。
- 分配的なグラフと表--ここで記述された確率分布ツールであり、同じ分布の異なったパラメーターを比較する為に使用されます(例、[2, 2], [3, 5]と[3.5, 8]のアルファとベータを持ったワイブル分布の形状とPDF、CDFとICDFの値と相互的なオーバーレイ)。
- オーバーレイグラフ--異なった分布を比較する為に使用され(理論的な入力仮定と実践的にシミュレーションされた予測結果)、表示的な比較のために相互的なオーバーレイをします。

過程:

- ROV BizStats を **リスクシミュレーター/ 分配的グラフと表** から起動し、**グローバルな入力の適用** のボタンをクリックして、入力パラメーターのサンプル設定を開くか、独自の入力値を記入して**実行**を選択します。生成された 4 つのモーメントと CDF、ICDF、PDF は、45 の分布の各 1 つずつに計算されます (Figure 5.48)。

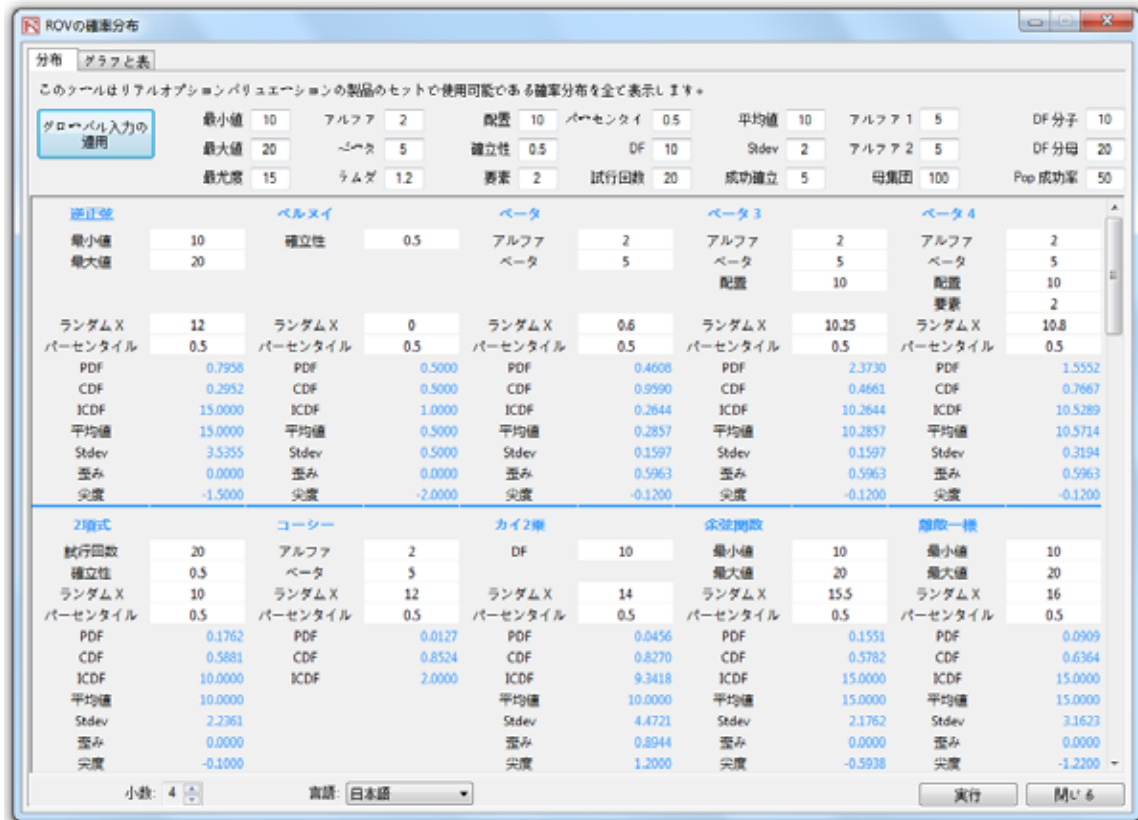


Figure 5.48 - 確率分布ツール (45 通りの確率分布)

- グラフと表** のタブ (Figure 5.49) をクリックし、分布を選択 **[A]** (例、逆正弦)。CDF、ICDF、あるいは PDF の実行を選択 **[B]**。重要な入力を記入し、**グラフの実行**、あるいは **表の実行** を選択 **[C]**。結果の表示にグラフと表のタブを帰ることができ、いくつかのグラフアイコンを試して、グラフへの変化を見ることがもできます **[E]**。
- 2 つのパラメーターの変換 **[H]** し、**～から/～まで/ステップ** 入力、あるいは **カスタム** 入力を選択し、**実行** をクリックすることで複数のグラフと表、そして分布表の生成が行えます。例えば、5.50 で表示されているように、ベータ分布を実行してから PDF を選択し **[G]**、カスタム入力 **[I]** の使用によって、交換するためのアルファとベータを選択し **[H]**、重要な入力パラメーターを記入します：2;5;5 はアルファ、5;3;5 はベータ **[J]**、そして **グラフの実行** をクリックします。これは、3 つのベータ分布を生成します **[K]**：ベータ (2, 5)、ベータ (5, 3) とベータ (5, 5) **[L]**。様々なグラフタイプ、グリッドライン、言語と小数を試してください **[M]**。また、理論 対

実践的にシミュレーションされた値を使用した分布の再実行も試してみてください[N]。

Figure 5.51 は、2 項式分布のために生成された確立表を表示しており、成功確率と成功試行回数(乱数 X)は **〜から〜まで/ステップ** オプションの使用を通して異なるように選択されています[0]。表示されている様に計算の複製を試して、表タブをクリック[P]し、作成された確立密度関数結果を確認します。この例証では、2 項式分布で、試行回数の初期の入力設定= 20、成功確率= 0.5 と成功試行回数 $X = 10$ に設定され、変数の行に表示される成功確率は 0. から 0.25、…、0.50 まで異なることを可能とし、変数の列に表示される成功試行回数も 0 から 1、2、…、8 まで異なることを可能とします。PDF の選択によって表の結果は与えられたイベント発生の確率を表示します。例えば、各試行回数の成功率が 25% の場合の 20 回の試行から 2 回の成功確立を得た場合、0.0669 、あるいは 6.69% の確立と表示されます。

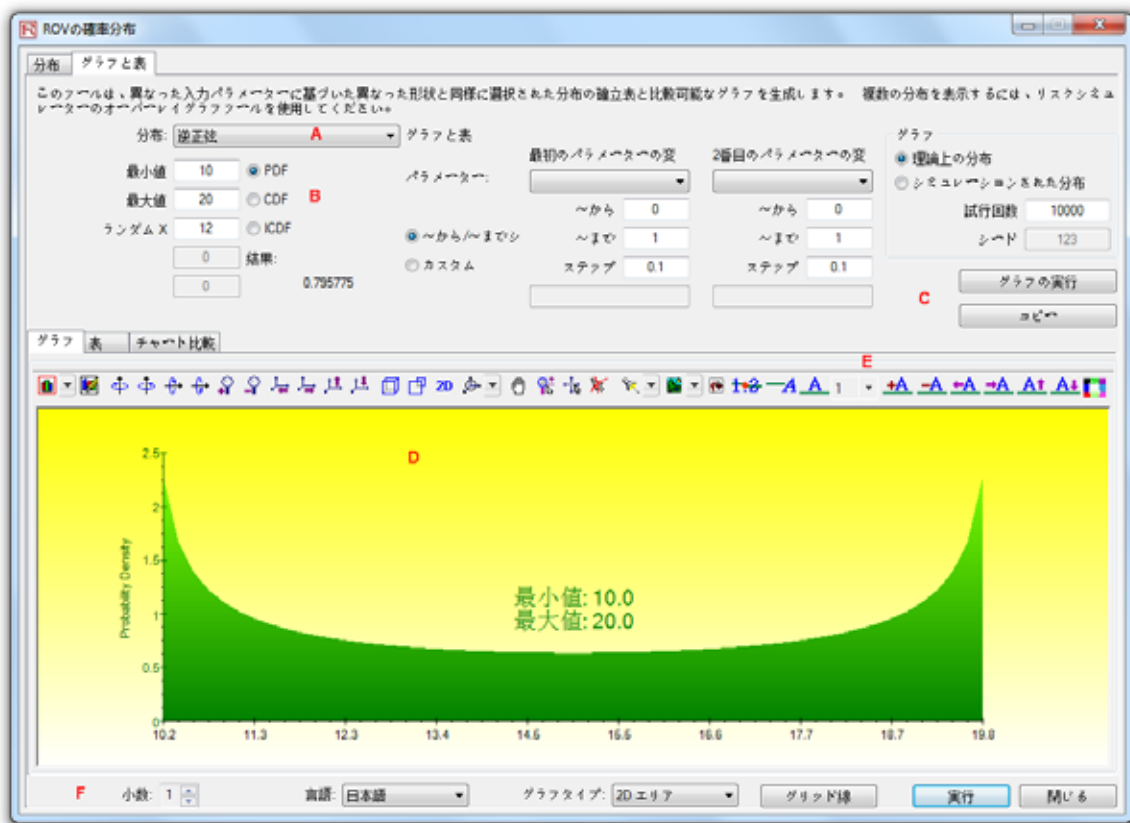


Figure 5.49 - ROV 確率分布(PDF と CDF グラフ)

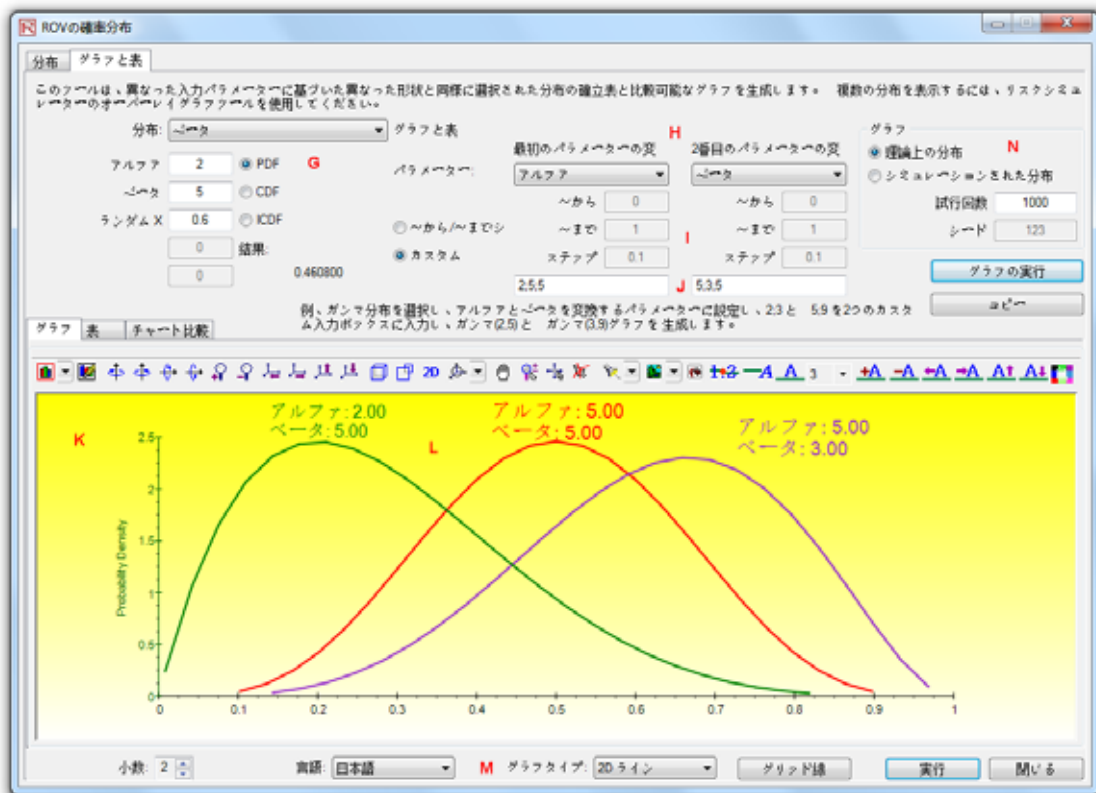


Figure 5.50 - ROV 確率分布 (複数のオーバーレイグラフ)

N	0.0000	1.0000	2.0000	3.0000	4.0000	5.0000	6.0000	7.0000	8.0000
1	0.2000	0.0115	0.0576	0.1369	0.2054	0.2182	0.1746	0.1091	0.0545
2	0.2500	0.0032	0.0211	0.0669	0.1339	0.1897	0.2023	0.1686	0.1124
3	0.3000	0.0008	0.0068	0.0278	0.0716	0.1304	0.1789	0.1916	0.1643
4	0.3500	0.0002	0.0020	0.0100	0.0323	0.0738	0.1272	0.1712	0.1844
5	0.4000	0.0000	0.0005	0.0031	0.0123	0.0350	0.0746	0.1244	0.1659
6	0.4500	0.0000	0.0001	0.0008	0.0040	0.0139	0.0365	0.0746	0.1221
7	0.5000	0.0000	0.0000	0.0002	0.0011	0.0046	0.0148	0.0370	0.0739

Figure 5.51 - ROV 確率分布 (分布表)

ROV BizStats

この新しい ROV BizStats ツールは、リスクシミュレーターのとて強力で速いモジュールであり、ビジネス統計やデーターの分析モデルを実行するために使用され、130 以上のビジネス統計と分析モデル (Figures 5.52-5.55) が含まれています。次にモジュールの実行においての簡単な導入とソフトウェアの各要素の詳細が記述されています。

過程:

- ROV BizStats は **リスクシミュレーター/ ROV BizStats** から起動することができ、**例証**をクリックするとサンプルデーターとモデルのプロファイルを開く **[A]** ことができます。また、データーの入力、あるいはコピー/貼り付けをデーターグリッド内で行うこともできます **[D]** (Figure 5.52)。最初のメモの行で独自のメモや変数名称を追加することができます **[C]**。
- 重要なモデルを選択 **[F]** し、例証データーの入力設定 **[G]** を使用して重要な変数への入力を行い **[H]** ステップ 2 に続行します。同じパラメーターには、セミコロンと新しい線を使用して (Enter をヒットして新しい線を作成) パラメーターの区分を行うことができます。
- 結果の計算 **[J]** には **実行** **[I]** をクリック。重要な分析結果、グラフ、あるいはステップ 3 からの様々な表の統計が表示されます。
- 必要である場合、保存のためにステップ 4 でプロファイル内にモデル名称を入力することが可能 **[L]** です。複数のモデルを同じプロファイル内に保存することができます。既存するモデルの編集、あるいは削除 **[M]** ができ、表示順に整理 **[N]** しなおすこともでき、全ての変更は、単一プロファイルに保存 **[O]** することがファイル名称に *.bizstats の拡張を付けることで可能となります。

メモ:

- データーのグリッドサイズは、メニューで設定でき、グリッドは、1,000 以上の変数列と変数につき百万行のデーターを調整することができます。また、メニューでは、言語設定とデーターの小数設定の変更もできます。
- はじめるにあたって、事前に作成された欠陥のないデーターといくつかのモデル **[S]** が含まれているため、例証ファイルを開く **[A]** ことをお勧めします。これらのモデルのどれかを選んでダブルクリックすることで実行し、レポートエリアに結果 **[J]** が表示されます。例証によってグラフ、あるいはモデル統計が表示されます **[T/U]**。この例証の使用によってモデルの記述 **[G]** に基づいた入力パラメーター **[H]** の記入されたかを参照することができ、独自のカスタムモデルを作成する過程に進むことができます。

- 変数のヘッダー[D]をクリックすることで、一度に 1 つあるいは複数の変数を選択することができ、右クリックで選択された変数の追加、削除、コピー、貼り付け、または表示[P]ができます。
- コマンド・コンソール[V/W/X]を使用してモデルを記入することができます。この作動を見るには、モデルの実行をダブルクリックし[S]、コマンド・コンソールに入ります[V]。モデルの複製、あるいは独自のモデルを作成し、コマンドの実行[X]をクリックします。コンソールでの各線はモデルと重要なパラメーターを表しています。
- *.bizstats プロファイル（データーと複数のモデルが作成され、保存された）全体は、ファイルメニューから XML の編集を開くことで直接 XML[Z]で編集することができます。プロファイルの変更は、プログラムのここで行うことができますが、保存する必要があります。
- グリッドの列ヘッダーをクリックして、列全体、あるいは変数の全てを選択し、ヘッダー上を右クリックをヒットすると列のデーターの自動適合、カット、コピー、削除、あるいは、貼り付けができます。また、複数の列ヘッダーをクリックして複数の変数を選択することができ、右クリックで表示するを選択するとデーターをグラフ化します。
- セルが大きな値を持っている場合、全ては表示されないため、クリックしたままマウスを値が表示されているセル上に持って行くと全ての値が表示されたコメント画面がポップアップされます。あるいは単に変数の列のサイズを変える事もできます(列の幅を広げ、列の枠をダブルクリックすると自動的に適合します。あるいは、列ヘッダー上を右クリックし、自動適合を選択します)。
- グリッド上を移動するのに上下、左右のキーの使用、あるいはキーボードのホームとエンドを使用して行の一番左側か一番右側に移動することが出来ます。また、次のような混合キーを使用する事も出来ます：Ctrl+ホームは、一番左のセルに移動、Ctrl+エンドは、右下のセル、シフト+上/下は、特定のエリアを選択するため等。
- メモの欄には、各変数の短いメモを記入することができます。短くて明確なメモを入力することを忘れないで下さい。
- グラフの見た目や読解さを工夫する為に表示するの選択で様々なグラフのアイコンを試みて下さい(例、回転、シフト、ズーム、色の変更、伝言の追加など)。
- コピーボタンは、結果、グラフと統計の表をモデルが実行された後のステップ 3 でコピーする為に使用されます。モデルが実行されていない場合は、空白のページをコピーすることになります。
- レポートボタンは、ステップ 4 でモデルが保存された場合、あるいはグリッドにデーターがある場合にだけ実行され、他の生成されるレポートは空白となります。また、Microsoft Excel のインストールが事前に必要とされていないと、データー抽出と結果レポートの実行できず、Microsoft PowerPoint がないとグラフのレポートが表示できません。
- 特定のモデル、あるいは統計手法の実行方法に疑問を持っている場合は、例証プロファイルを開き、ステップ 1 でどのようにデーターが設定されているかを確認、またはステップ 2 でどのように入力パラメーターが設定されているかを確認してください。最初はこれらの例証を独自のモデルの作成ガイドとして使用することができます。

- 言語は、**言語**メニューで変換することが出来ます。現時点では、10 言語の選択が出来る事が出来ます。ただし、時々、特定の有限結果は英語で表示され続けます。
- **ステップ 2** では、モデルの表示順位を **ビュー** で変更することが出来ます。モデルをアルファベット順、カテゴリー順、データー入力の必要条件順などに並べ替えることが出来ます。特定のユニコード言語では (例、中国語、日本語、韓国語など)、アルファベット順が適応しない為、その選択肢は無効となります。
- ソフトウェアは、様々な地域の小数と数字の設定を操作することが出来ます(例、千ドル五十セントは 1,000.50、1.000,50、あるいは 1'000,50 と表示することが出来ます)。小数の設定は、ROV BizStats のメニューの **データー | 小数の設定** で指定することが出来ます。どっちにしろ、疑問がある場合は、ROV BizStats でコンピューターの地域設定を米国英語にし、デフォルトである北米式の 1,000.50 に設定してください(この設定は、ROV BizStats やデフォルトの例証の確実な操作を保証します)。

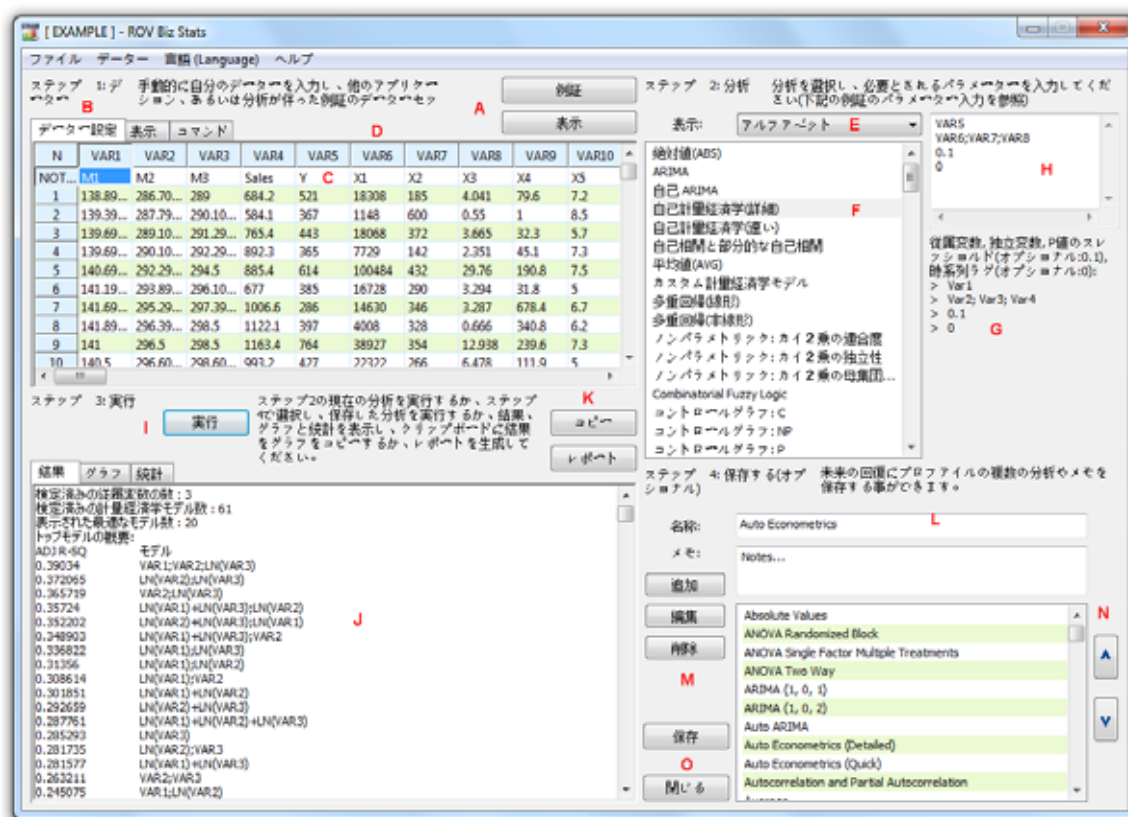


Figure 5.52 - ROV BizStats (統計分析)



Figure 5.53 - ROV BizStats (データの表示と結果グラフ)

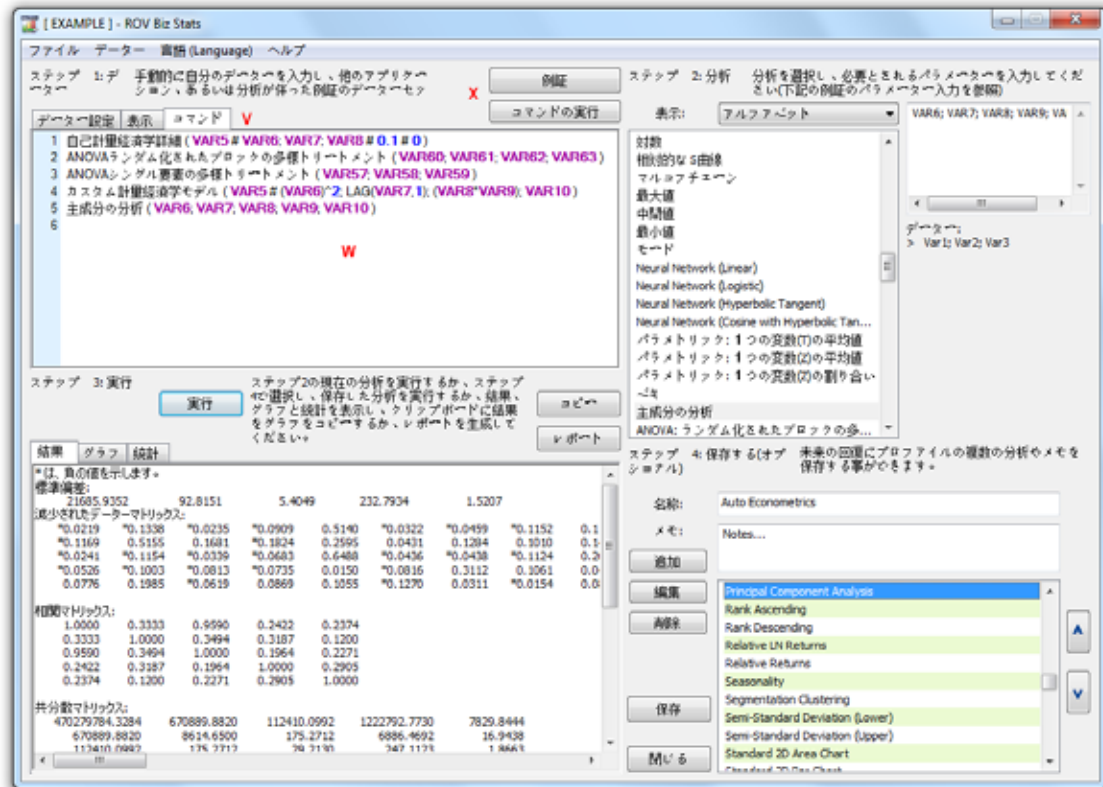


Figure 5.54 - ROV BizStats (コマンド・コンソール)

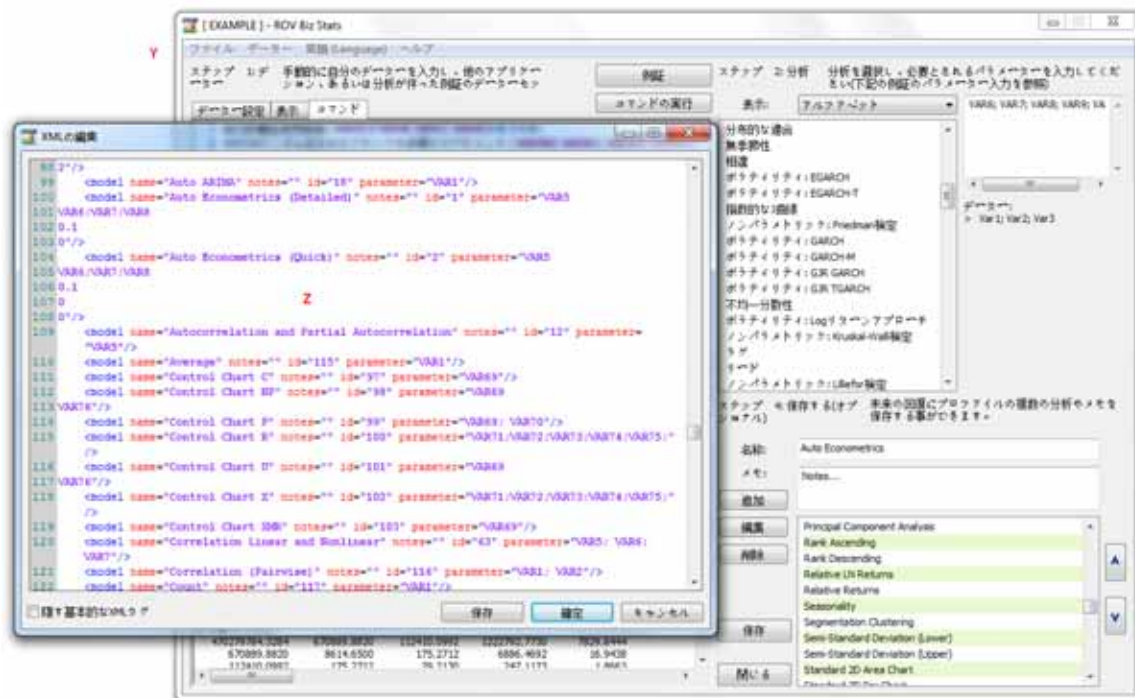


Figure 5.55 - ROV BizStats (XML の編集)

ニューラル・ネットワークと組み合わせ的ファジィ論理予測法

ニューラル・ネットワークの概念は、たいてい生物学的な神経細胞の回路やネットワークを示すのに使用されますが、現代的な使用方法として、ソフトウェアの環境で再現された人工的なニューラル、あるいはノードで構成された人工的なニューラル・ネットワークを示します。この方法は、人間の脳、あるいはニューロンの考え方、パターンの識別、我々のケースでは、予測法の時系列データーの目的の為のパターンの識別などを複製することを試みています。この手法は、**ROV BizStats** モジュールのリスクシミュレーター内にあり、細かくは、[リスクシミュレーター| ROV BizStats | ニューラル・ネットワーク](#)から、または[リスクシミュレーター| 予測法 | ニューラル・ネットワーク](#)からアクセスすることができます。Figure 5.56 は、ニューラル・ネットワークの予測法を表示しています。

過程

- [リスクシミュレーター| 予測法 | ニューラル・ネットワーク](#)をクリック。
- 手動的にデーターを入力するか、クリップボードからあるデーターを貼り付け(例、Excel からあるデーターを選択し、コピーして、このツールを起動してデーターを貼り付けます)るかのどちらかを実行。
- [線形](#)、あるいは [非線形](#)ニューラル・ネットワークモデルのどちらかを選択し、望む [予測周期数](#) (例、5)、ニューラル・ネットワーク(例、3)の隠れた [レイヤー数](#) と [検定周期数](#) (例、5)を入力。
- [実行](#)をクリックして、[分析](#)をし、計算された結果とグラフを確認。また、結果とグラフをクリップボードに [コピー](#)し、他のソフトウェア・アプリケーションに貼り付けることが可能。

ワークブックの隠れたレイヤー数は、入力パラメーターであり、データーと調節する必要があることに注意してください。一般的に、複雑なデーターパターンであるほど、多数の隠されたレイヤーが必要となり、計算にもより時間が掛かります。3 つのレイヤーから始めることをお勧めします。検定周期は、ニューラル・ネットワークモデルの最終調節で使用されたデーターポイント数であり、少なくとも予測周期数と検定周期数が同じであることを勧めます。

一方、ファジィ論理の概念はファジィ論理の設定から成り立ったもので、普通の論理の反対である正確よりもおよそである論理を取り扱い、2 値の設定は 2 値の論理を持ち、ファジィ論理の変数は、0 と 1 の間に相当する値があり、典型的な論理の 2 つの値の規制などはありません。このファジィなスキーマは、組み合わせ的な手法と一緒に使用され、Figure 5.57 で表示されているようにリスクシミュレーターの時系列予測結果を与え、季節性と傾向性を持っている時系列データーへの適用はさらに有効です。この手法は、

ROV BizStats モジュールのリスクシミュレーター内の[リスクシミュレーター | ROV BizStats | 組み合わせ的なファジィ論理](#)から、あるいは[リスクシミュレーター | 予測法 | 組み合わせ的なファジィ論理](#)アクセスすることができます。Figure 5.57 は、ニューラル・ネットワーク予測法を表示しています。

過程

- [リスクシミュレーター | 予測法 | 組み合わせ的なファジィ論理](#)をクリック。
- 手動的にデータを入力するか、クリップボードからあるデータを貼り付け (例、Excel からあるデータを選択し、コピーして、このツールを起動してデータを貼り付けます) するかのどちらかを実行。
- 一覧表から実行したい分析の変数を選択し、季節性周期と (例、4 は、4 半期データに、12 は、月間データになど)、望む [予測周期](#) を入力します (例、5)。
- [実行](#) をクリックして、分析をし、計算された結果とグラフを確認。また、結果とグラフをクリップボードに [コピー](#) し、他のソフトウェア・アプリケーションに貼り付けることが可能。

ニューラル・ネットワークとファジィ論理技法のどちらもビジネス予測法ドメイン、および戦略、攻略、または作業レベルで有効、および信頼性のある手法として定義されていないことに注意してください。これらの高度な予測法のフィールドでは、沢山の研究や課題があるにも関わらず、リスクシミュレーターは、これらの 2 つの技法の基本を提供し、時系列予測法の実行目的をサポートしています。孤立でこれらの技法を使用することをお勧めしませんが、他のリスクシミュレーターの予測技法と共により強硬なモデルを構築するに組み合わせるのは、いい考えでしょう。

ニューラル・ネットワーク予測法

ステップ 1: データを手動的に自分のデータベースを入力し、他のアプリケーション、あるいは分析が伴った例証のデータベースを開いてください。

貼り付け

N	VAR1	VAR2	VAR3	VAR4	VAR5	VAR6	VAR7	VAR8	VAR9	VAR10
NOT...	NNET									
1	459.11									
2	460.71									
3	460.34									
4	460.68									
5	460.83									
6	461.68									
7	461.66									
8	461.64									
9	465.97									
10	469.38									

ステップ 2: 分析タイプ、変数と実行する予測周期の選択。

日本語

VAR1

3

210

210

コピー

実行

結果 グラフ

☒ 複数段階の最適化を適用

Sum of Squared Errors (Training) : 1.745466
 RMSE (Training) : 0.091827
 Sum of Squared Errors (Modified) : 79550.322561
 RMSE (Modified) : 19.463069
 Forecasting
 * indicates negative values

Period	Actual (Y)	Forecast (F)	Error (E)
211	581.5000	616.2004	*34.7004
212	584.2200	616.8210	*32.6010
213	589.7200	617.4589	*27.7389
214	590.5700	617.9290	*27.3590
215	588.4600	618.4628	*30.0028
216	586.3200	619.4553	*33.1353
217	591.7100	619.7934	*28.0834
218	593.2600	619.9928	*26.7328
219	592.7200	620.6259	*27.9059
220	592.3000	621.3075	*29.0075
221	589.2900	621.8699	*32.5799
222	593.9600	622.1975	*28.2375
223	597.3400	622.7418	*25.4018
224	600.0700	622.9101	*22.8401
225	596.8500	622.6457	*25.7957

Figure 5.56 – ニューラル・ネットワーク予測法

組み合わせのファジィ論理予測法

ステップ 1: データを自動的に自分のデータを入力し、他のアプリケーション、あるいは分析が伴った例証のデータセットを開いてください。 貼り付け

N	VAR1	VAR2	VAR3	VAR4	VAR5	VAR6	VAR7	VAR8	VAR9	VAR10
NOT...		FUZZY								
1	459.11	684.2								
2	460.71	584.1								
3	460.34	765.4								
4	460.68	892.3								
5	460.83	885.4								
6	461.68	677								
7	461.66	1006.6								
8	461.64	1122.1								
9	465.97	1163.4								
10	469.38	993.2								

ステップ 2: 必要な入力を入力し、予測の目的変数を選択。

日本語

VAR2

季節性: 4

予測周期: 10

コピー 実行

結果 グラフ

Results RMSE : 707.039492
 Auto ARIMA RMSE : 249.495091
 Time-Series Auto RMSE : 287.252763
 Trend Line Exponential RMSE : 775.403678
 Trend Line Linear RMSE : 912.616213
 Trend Line Logarithmic RMSE : 1488.012692
 Trend Line Moving Average RMSE : 988.333906
 Trend Line Polynomial RMSE : 758.307610
 Trend Line Power RMSE : 1268.660480

RESULTS

Forecast Fit

* indicates negative values

Period	Actual (Y)	Forecast (F)	Error (E)
1	684.2000		
2	584.1000		
3	765.4000		
4	892.3000		
5	885.4000	802.4484	82.9516
6	677.0000	863.9179	*186.9179
7	1006.6000	971.7020	34.8980
8	1122.1000	1083.6028	38.4972

Figure 5.57 – ファジィ論理時系列予測法

ゴールシークの最適化

ゴールシーク・ツールは、適用されるアルゴリズム的な検索であり、モデル内の単一変数の解決を見出します。公式、あるいはモデルから得たい結果を知っているが、その結果を得るためにどの入力値を記入しなければならないのかが明確でない場合、[リスクシミュレーター | ツール | ゴールシーク](#) を使用してください。ゴールシークは、単一変数値でしか実行できないことに注意してください。1 つ以上の入力値を認知したい場合、リスクシミュレーターの高度な最適化ルーチンを使用してください。Figure 5.58 は、サンプルなモデルとゴールシークの適用を表示しています。

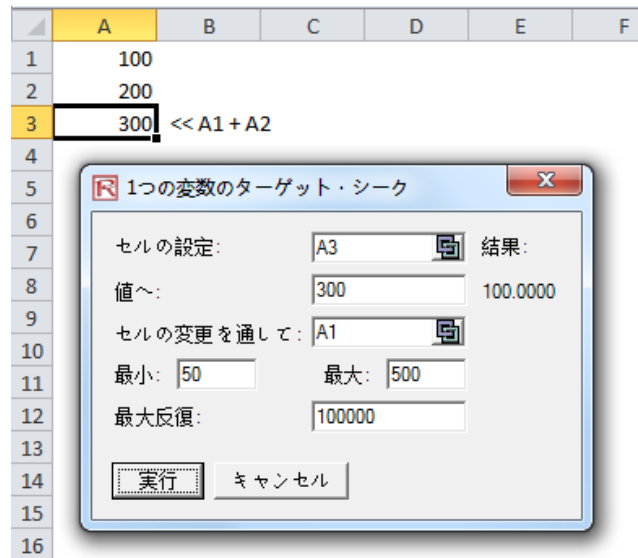


Figure 5.58 –ゴールシーク

単一変数の最適化

単一変数の最適化ツールは、モデル内の単一変数の解決を見出すために使用されるアルゴリズム的な検索であり、先例で記述されたゴールシーク・ルーチンと全く同じです。モデルからできる限りの最大値、あるいは最小値の結果を得たいが、その結果を得るためにどの入力値を記入しなければいけないのかが明確でない場合、**リスクシミュレーター | ツール | 単一変数の最適化**を使用してください(Figure 5.59)。この単一変数の最適化の実行はとても速いが 1 つの変数入力しか見出さないことに注意してください。1 つ以上の入力値を認知したい場合、リスクシミュレーターの高度な最適化ルーチンを使用してください。このツールがわりやすくシミュレーターに含まれているのは、時々、単一決定変数の素早い最適化の計算を必要とし、このツールは、プロファイル、シミュレーション仮定、決定変数、目的と規制などを含めた最適化モデルの設定なしで実行する性能を持っています。

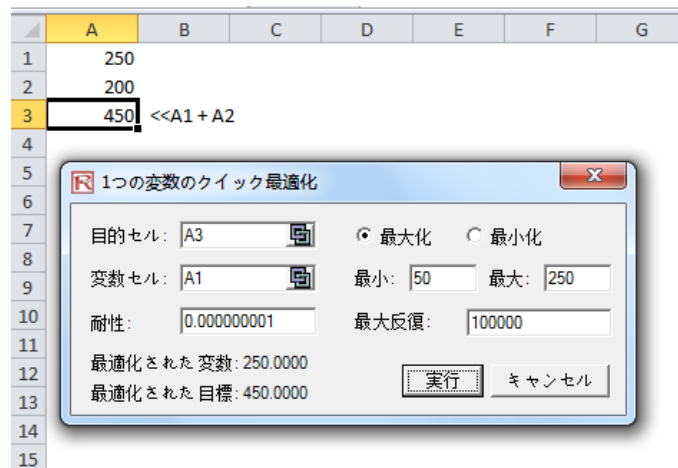


Figure 5.59 – 単一変数の最適化

遺伝子的アルゴリズムの最適化

遺伝子的アルゴリズムは、ヒューリスティック的な搜索であり、自然の進展の過程を再現しています。このヒューリスティックは、最適化と問題の搜索に便利な解決を生成するために使用されています。遺伝子的アルゴリズムは、進化のアルゴリズムの大きなクラスに属し、遺伝、突然異変、選択とクロスオーバーなどのような自然の進展に基づいた技法を使用して最適化問題の解決を生成します。

遺伝子的なアルゴリズムは、[リスクシミュレーター | ツール | 遺伝子的アルゴリズム](#) (Figure 5.60)からアクセスができます。結果は入力にとっても敏感なため(デフォルト入力は、最も多い入力レベルの一般的なガイドとして提供されます)、モデルの入力の調節に注意を払うことは必要です。結果のより強硬な設定の為にグラジェント搜索検定を選択することをお勧めします。(最初はこのオプションの選択を外し、その後このオプションを戻し、分析を再実行して結果を比較してみてください)。

メモ: 多くの問題では、遺伝子的アルゴリズムはローカル・オプティマのポイントか、問題のグローバル・オプティマよりも専制的なポイントのどちらかに集まる傾向を見せます。これは、短期的な適合を犠牲することでより長期的な適合を得るノウハウを持っていないことを示します。特定の最適化の問題と問題例には、遺伝子的アルゴリズムではない、より良い他の解決法が見つかるでしょう(同じ計算時間を与える)。従って、まずはじめに遺伝子的アルゴリズムを実行し、モデルの強硬性を確認する為に[グラジェント搜索検定の適合](#) (Figure 5.60)を選択してから再実行することをお勧めします。このグラジェント搜索検定は、典型的な最適化技法と遺伝子的アルゴリズム技法の組み合わせを試みた技法であり、出来る限り最的な解決を与えます。最後に、遺伝子的アルゴリズムを使用するには特定な論理的な必要条件がありますが、リスクシミュレーターの最適化モジュールの使用でより強硬な結果が得ることをお勧めします。ちなみに、より高度なリスクに基づいたダイナミック、およびストキャスティック最適化ルーチンの実行も可能にします。

ROV 決定木モジュール

決定木

Risk Simulator | Decision Tree 決定木モジュールを実行します(Figure 5.61).

ROV 決定木は、決定木モデルの作成と評価に使用されます。また、付加の高度な方法や分析が含まれています。

- 決定木モデル
- モンテカルロ・リスクシミュレーション
- 感度解析
- シナリオ分析
- ベイズ (ジョイントと事後確立の更新)
- 情報の予測値
- ミニマックス
- マクシミニ
- リスクプロファイル

決定木モジュールは、戦略木モジュールと同じ機能を持っていますが(詳細には、戦略木モジュールのセクションを参照)、次のような付加分析が含まれています。

- 既存するのどのノードのどれかを選択して、決定ノードのアイコン(四角)、不確実性ノードのアイコン(丸)、または、ターミナルノードのアイコン(三角)をクリックするか、メニューの挿入機能を使用して、決定、不確実性、あるいは、ターミナル・ノードの挿入することができます。
- 個々の決定、不確実性、あるいはターミナルノードののプロパティをノード上をダブルクリックして変更します。次に表示されているアイテムは、決定木モジュールの独特ないくつかの付加アイテムで、ユーザーインターフェースのノードのプロパティでカスタム化や設定が行えます。
 - 決定ノード: ノード上の値をカスタムオーバーライド、あるいは、自動的に計算します。自動計算のオプションは、デフォルトで設定されていて、完全な決定木モデルで実行をすると、決定ノードは結果と同時に更新されます。
 - 不確実性ノード: イベント名、確立やシミュレーション仮定の設定。確立イベント名、確立とシミュレーションの追加ができますが、既に不確実性のブランチが作成されていないといけません。
 - ターミナルノード: 手動入力、Excel リンクとシミュレーション仮定の設定。ターミナルイベントの支払いは、手動的に入力するか、Excel のセルにリンクする(例、支払いを計算する大規模な Excel モデルの場合、モデルをこの Excel モデルの結果セルにリンクすることができますあるいは、シミュレーション実行の為の確率分布仮定の設定をすることができます。
- ノードプロパティ画面の表示は、メニューの編集からアクセスでき、選択されたノードのプロパティは、ノードの選択と同時に更新されます。
- ちなみに決定木モジュールは、次の高度な分析ツールが含まれています。
 - 決定木上のモンテカルロ・シミュレーションのモデリング
 - 事後確立の取得の為のベイズ分析
 - 完全な情報の予測値、ミニマックス、マクシミニ解析、リスクプロファイルと不完全な情報の値

- 感度解析
- シナリオ分析
- ユーティリティ機能解析

シミュレーションのモデリング

このツールは、モンテカルロ・リスクシミュレーションを決定木上で実行します(Figure 5.62)。このツールは、シミュレーションの実行に確率分布を入力仮定として設定することを可能とします。また、選択したノードへの仮定、あるいは、新しい仮定を設定することもでき、これらの新しい仮定(あるいは、以前に作成した仮定)を公式などに使用することができます。例、正規という名の新しい仮定を(例、100の平均値と10の標準偏差を持った正規分布)を設定し、シミュレーションを決定木で行うか、この仮定を(100*正規+15.25)のような公式で使用できます。数値的ボックスに独自のモデルを作成します。基本的なシンプルな計算を使用するか、既存する変数の一覧表をダブルクリックして公式に既存する変数を追加することができます。数値でも、仮定でも新しい変数として一覧表に必要なだけ追加することができます。

ベイジアン分析

このベイジアン分析ツールは、パスによって関連されている2つの不確実なイベントを実行することができます(Figure 5.63)。例えば、右の例証では、不確実 A と B が関連されている年、タイムライン上でイベント A が先に発生し、イベント B がその後発生するとします。最初のイベント A は、2つの結果を持ったマーケットリサーチ(好都合か不都合)です。2番目のイベント B も 2つの結果を持ったマーケット条件(強い弱いか)です。このツールは、ジョイント、マージナルを計算する為に使用され、その後、ベイジアンは、優勢な確立と信頼性のある条件確立を入力することで確立を更新します。また、信頼確立は、条件確立を更新した後に計算することもできます。希望する重要な分析を下記から選択し、例証を開くをクリックして、選択された分析に対応するサンプル入力を参照すると、右のグリッドに結果の他にどの結果が決定木で入力として使用されているのかが表示されます。

- ステップ 1: 1 番目と 2 番目の不確実なイベントの名前を入力し、どれくらいの確立イベントを各自持っているかを選択します(自然状態、あるいは結果)
- ステップ 2: 各確率イベントの名前、あるいは結果を入力する
- ステップ 3: 2 番目のイベント前の角イベント、あるいは結果の確立と条件確立を入力する。確立はの合計は、100%でなければいけません。

完璧な情報、ミニマックス、マクシミニ分析、リスクプロファイル値と完璧ではない情報の値の予測

このツールは、完璧な情報(EVPI)、ミニマックス、マクシミニ分析の予測値の他、リスクプロファイルと完璧ではない情報の値を計算します(Figure 5.64)。はじめるにあたって、決定ブランチ数、あるいは、考慮している戦略を入力し(例、小・中・大施設の建設)、不確実イベント数、あるいは自然結果の状態を各パスの下で(例、好都合なマーケット、都合の悪いマーケット)入力し、予想支払いを各シナリオの下で入力します。

完璧な情報の (EVPI) の予測値は、例、完璧な先読みと段取りを持っている事を前提として(より確実な確立結果を得る為のマーケット・リサーチ、あるいは他の平均値を通して)、自然の確立状態のより繊細な推測に対して、EVPI は、その情報に追加された値があれば計算を行います。(例、マーケットリサーチが値を追加した場合)。はじめるにあたって、決定ブランチ数、あるいは、考慮している戦略を入力し(例、小・中・大施設の建設)、不確実イベント数、あるいは自然結果の状態を各パスの下で(例、好都合なマーケット、都合の悪いマーケット)入力し、予想支払いを各シナリオの下で入力します。

ミニマックス(最大回帰を最小化する)とマクシミニ(最小支払いを最大化する)は、2つの対立的なアプローチであり、最適な決定パスを見出します。これらの2つのアプローチは、余り使用されませんが、決定判断仮定に潜在的な考えを貢献しています。既存する決定ブランチ数、あるいはパスを入力し(例、小・中・大施設の建設)、不確実イベント数、あるいは自然結果の状態を各パスの下で(例、好都合な経済、都合の悪い経済)入力し、いくつかのシナリオの為に支払い表を埋め合わせ、ミニマックスとマクシミニの結果を計算します。例証の計算を例証を開くをクリックすることで参照することができます。

感度性

入力確立上の感度性分析は、決定パスの値への影響を定義する為に実行されます。まず始めに、次の分析する1つの決定ノードを選択し、1つの確立イベントを一覧表から選びます(Figure 5.65)。同様な確立を持った複数の不確実性イベントがある場合、それらを個々で、あるいは一緒に分析することができます。

感度性グラフは、確立レベルによる決定パスの値を表示します。数値的な値は結果表に表示されます。クロスオーバーラインの配置は、場合によって、どの確立イベントに特定の決定パスが他よりも優勢なのかをあらわします

シナリオ表

シナリオ表は、入力にいくつかの変更が与えられた結果を定義する為に生成されます。分析する1つ、あるいはそれ以上の決定パスと(選択された各パスの結果は、別々の表とグラフに表示されます)、1つ、あるいは2つの不確実性、あるいはターミナルノードは、シナリオ表の入力変数として選択することができます(Figure 5.66)。

- 次の一覧表から分析する1つ、あるいはそれ以上の決定パスを選択する。
- モデルにする1つ、あるいは2つの不確実性イベント、あるいは支払いターミナルを選択する
- イベントの確立を独自に変更するか、同じような確立イベントのすべてを一度に変更するかを決めてください。
- <入力シナリオ範囲を入力。

ユーティリティ機能の生成

ユーティリティ機能、または、 $U(x)$ は、時々決定木の支払いターミナルの予測値の代わりに使用されていることがあります。 $U(x)$ は、2通りで向上することができます(Figure 5.67):

すべての確立結果の厳密で細かい実験を通してか、あるいは、指数的外押法(個々で使用されている)。

それらは、リスクを嫌がる決定判断者(下部は、上部より損害が多く、悲惨)、リスク・ニュートラル(上部と下部は誘引効果を持っている)、あるいはリスク好き(上部がより誘引効果を持っている)の為にモデル化することができます。支払いターミナルの最小と最大予測値とその相田のデータポイント数を入力し、ユーティリティ曲線と表を計算します。

50:50のギャンブルをしている場合、 $\$X$ を得るか、 $-\$X/2$ を負けるかに対して、全然賭けないで $\$0$ 得た場合、 $\$X$ の数値はどのくらいでなければいけないのか?例えば、 $\$100$ を得ることができる、あるいは $-\$50$ 負けるかもしれない賭け事に興味が無い場合、 X は、 $\$100$ となります。 X を下のポジティブな利益ボックスに入力します。 X の数値が大きければ、リスク好きではないことを示し、 X が小さければリスク好きであることを示します。

必要な入力を記入し。U(x)タイプを選択し、結果を得る為にユーティリティを計算するをクリックしてください。これを再実行するのに計算された U(x)の値を決定木に適用することができます。また、木を戻して、支払いの予測値を使用することもできます。

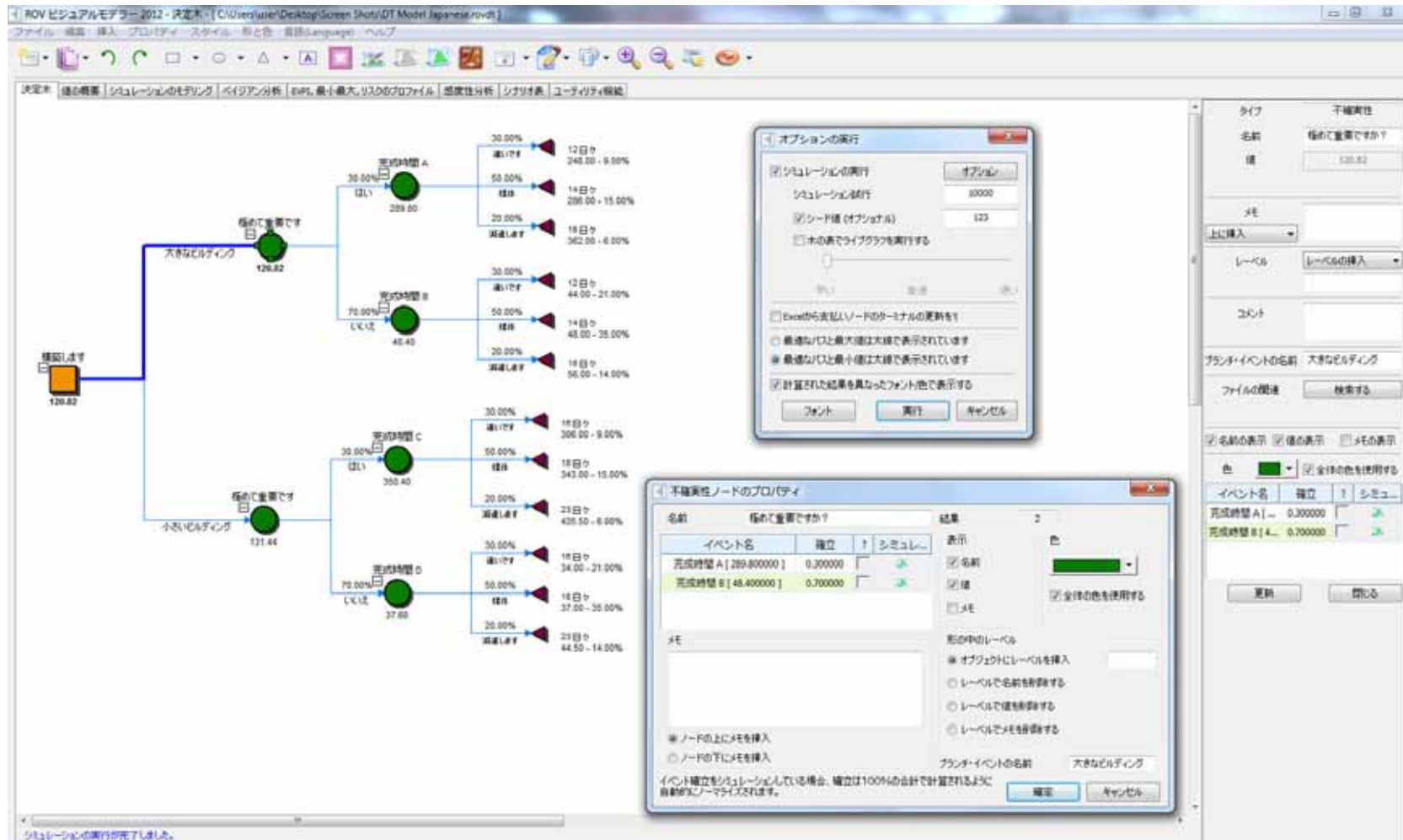


Figure 5.61 – ROV 決定木 (決定木)

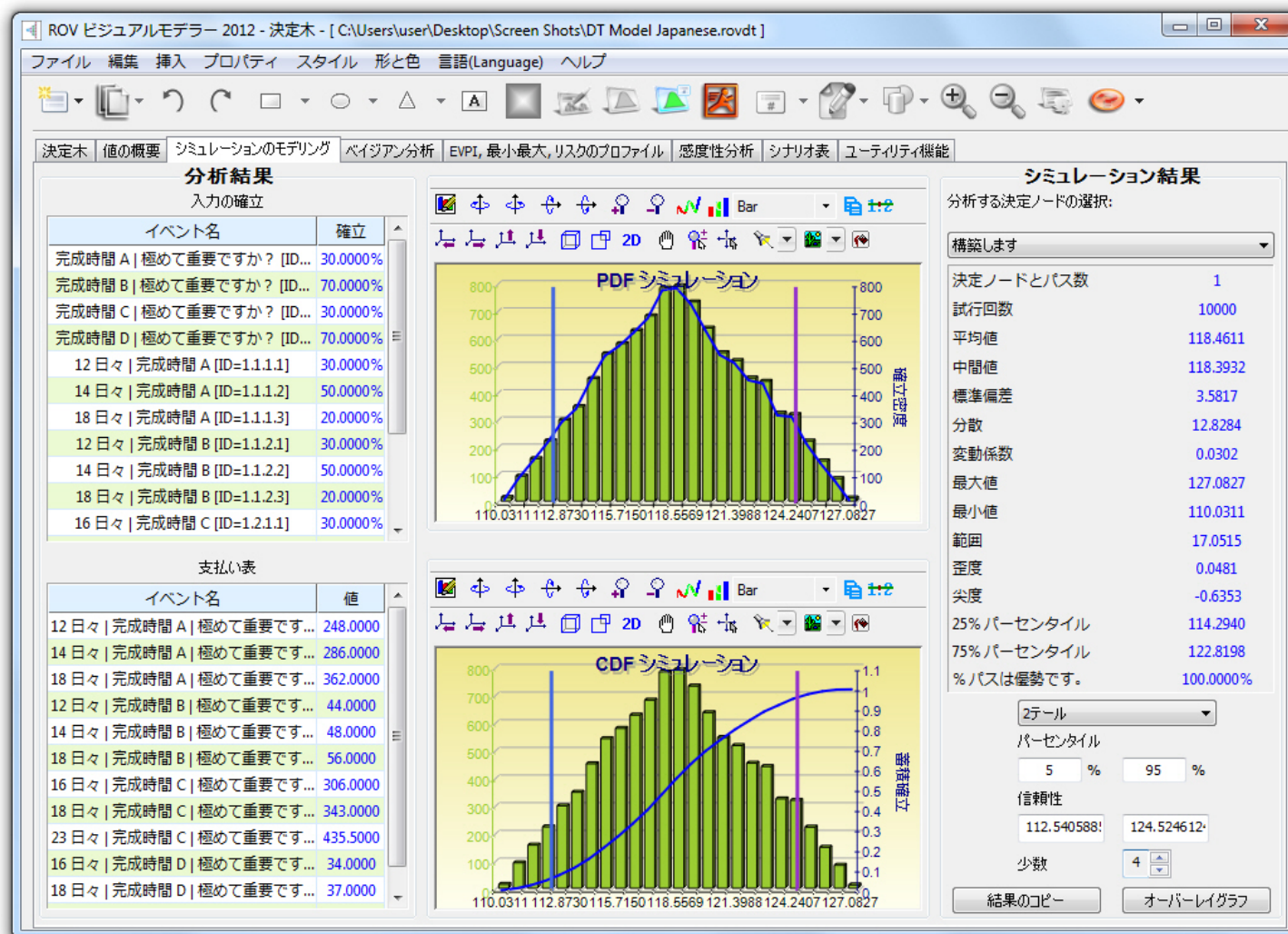


Figure 5.62 – ROV 決定木(シミュレーション結果)

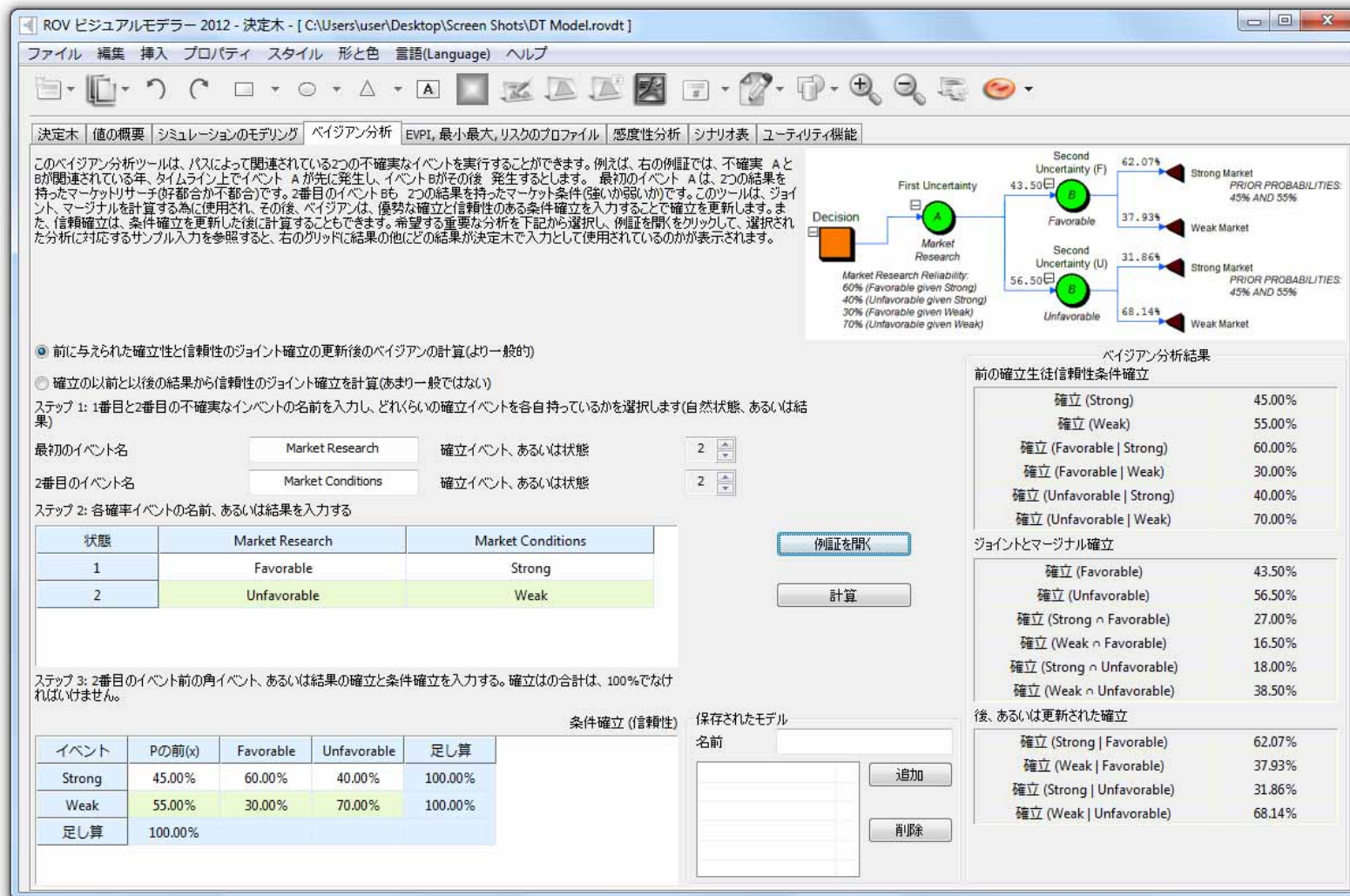


Figure 5.63 – ROV 決定木 (ベイズ解析)



Figure 5.64 – ROV 決定木 (EVPI, MINIMAX, リスクプロファイル)

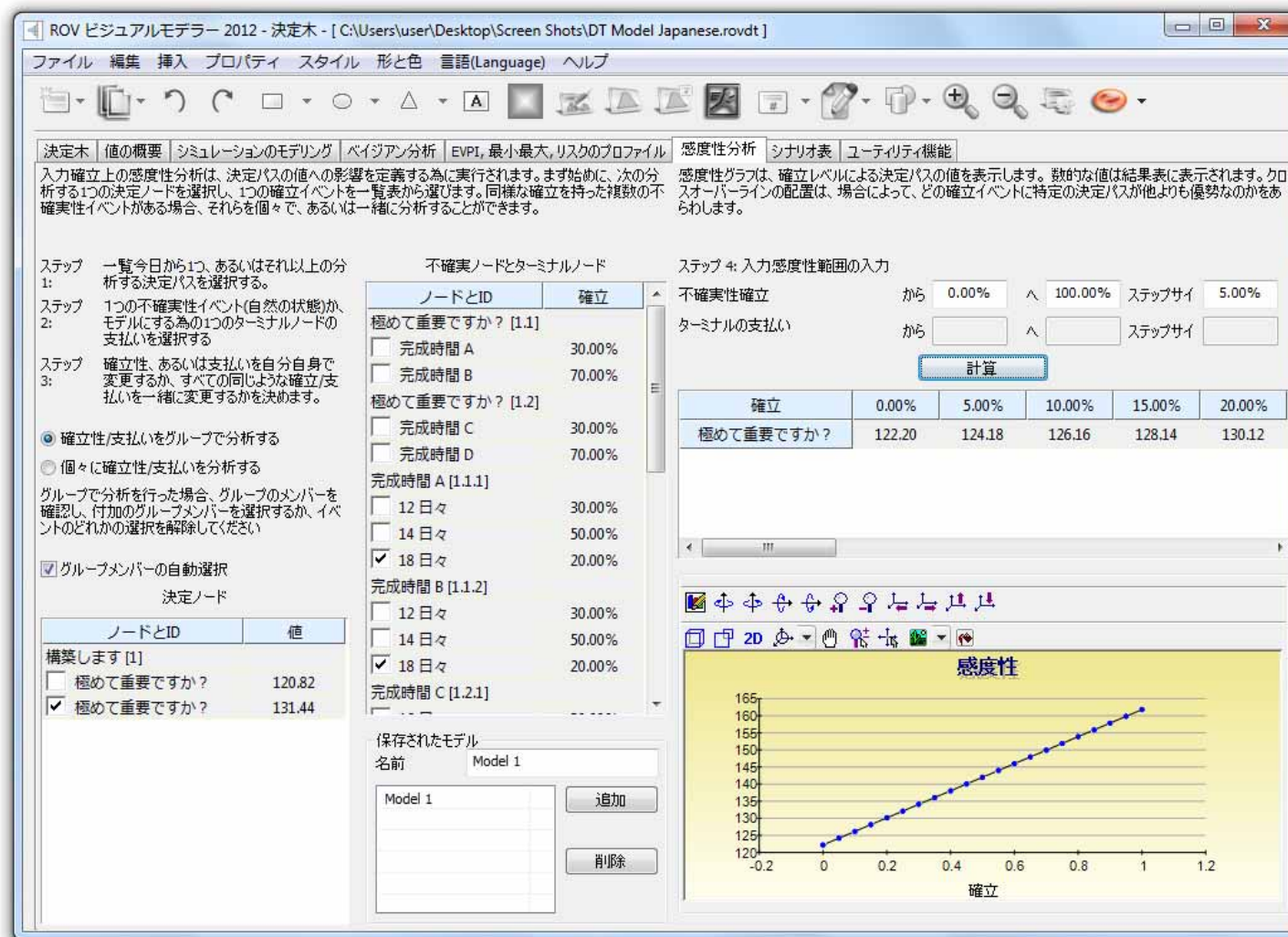


Figure 5.65 – ROV 決定木(感度解析)

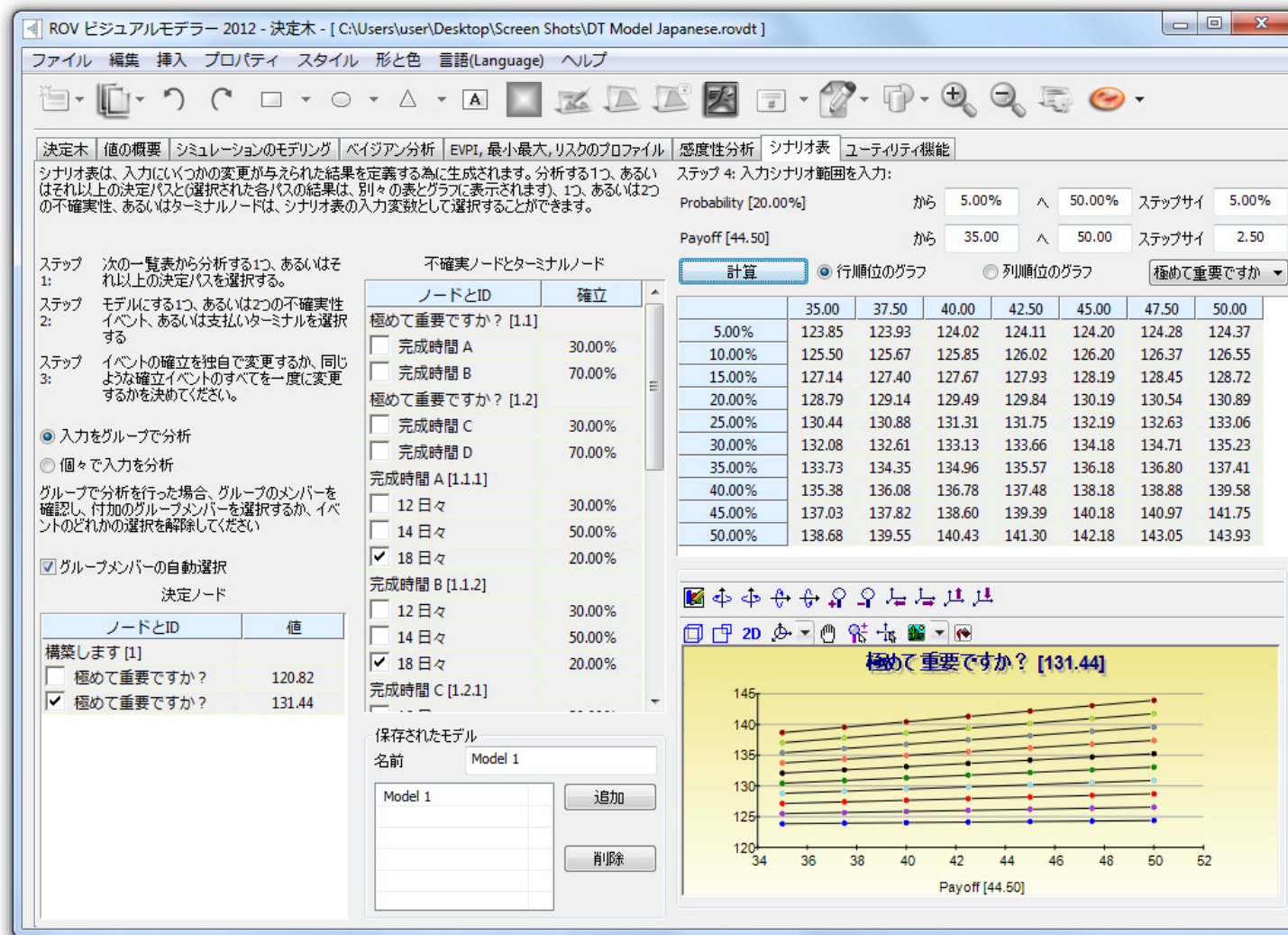


Figure 5.66 – ROV 決定木 (シナリオ表)

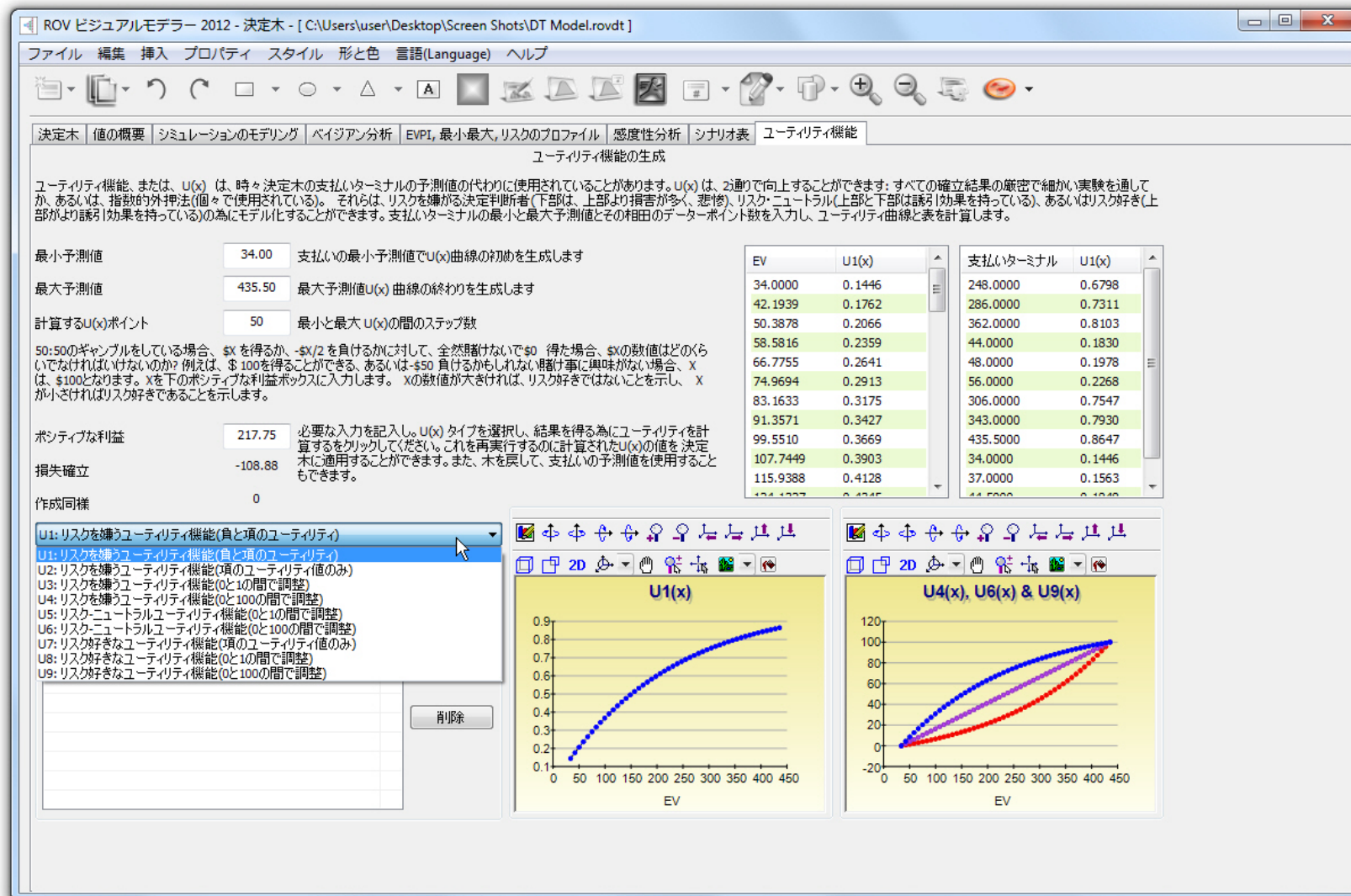


Figure 5.67 – ROV 決定木 (ユーティリティ機能)

役立つ情報と技法

次に高度なリスクシミュレーターユーザーに役立つ情報と簡単な技法が記述されています。特定のツールの詳細には、ユーザーマニュアルのあたいする章をご覧ください。

TIPS: 仮定(入力仮定ユーザーインターフェースの設定)

- クイック・ジャンプ--分布のどれかを選択し、ある文字を打つするとその文字で始まる最初の分布を表示します(例、正規をクリックし、wを打つとワイブル分布へ移動します)。
- 右クリック表示--分布のどれかを選択し、右クリックで分布の異なった表示を選択することができます(大きいアイコン、小さいアイコン、一覧表)。
- グラフ更新のためのタブ--新しい入力パラメーターの入力後(例、新しい平均値、あるいは標準偏差値の入力)、キーボードのタブを打つと、あるいは入力ボックス上とは違うユーザーインターフェース上のどこかをクリックすることで分布グラフが自動的に更新されます。
- 相関の入力--ペアワイズ相関をここから直接的な記入(列のサイズは必要に応じて変えることができます)と、複数の分配的な適合ツールの使用で全てのペアワイズ相関の記入と計算を行い、あるいは、いくつかの仮定の設定後、相関ツールの編集を使用して相関マトリックスを入力することができます。
- 仮定セルの公式--空欄セル、あるいは静的値を持ったセルだけが仮定として設定できますが、仮定セルで公式、あるいは関数が必要とされる場合、時間がなければいけません。この過程として、まずは入力仮定をセルに記入し、公式、あるいは関数(シミュレーション実行の際に、シミュレーションされた値が関数を置き換え、シミュレーションの完了後、関数、あるいは公式がまた表示されます)を記入します。

TIPS: コピーと貼り付け

- エスケープを使用したコピーと貼り付け--セルを選択し、リスクシミュレーターのコピー機能の使用すると、セルの値、公式、関数、色、文字とサイズその他、リスクシミュレーターの仮定、予測結果、決定変数などの全てがWindowsのクリップボードにコピーされます。次に貼り付けの過程として、2つの選択があります。最初の選択は、リスクシミュレーターの貼り付けを直接行い、全てのセルの値、色、文字、公式、関数とパラメーターを新しいセルに貼り付ける手段です。2つ目の選択として、キー

ボードでエスケープを打った後にリスクシミュレーターの貼り付けを行うことです。エスケープは、リスクシミュレーターにリスクシミュレーターの仮定、予測結果、決定変数だけを貼り付けたいと命じ、セルの値、色、公式や関数などは貼り付けません。貼り付けを行い前にエスケープを打つことは、目的セルの値と計算を維持することができ、リスクシミュレーターのパラメーターだけを貼り付けます。

- 複数のセルのコピーと貼り付け--コピーと貼り付けに複数のセルの選択が可能です(連続的な、そして連続的でない仮定)。

TIPS: 相関

- 仮定の設定--入力仮定ダイアログの設定を使用してペアワイズ相関の設定を行います(いくつかの相関だけの記入に理想的)。
- 相関の編集--手動的に記入、あるいはクリップボードの貼り付けによって相関マトリックスの設定を行います(大きな相関マトリックスと複数の相関に理想的)。
- 複数の分配的な適合--自動的にペアワイズ相関を計算し、記入します(複数の変数の適合を実行する際に理想的で、自動的に相関を計算し、統計的に有意な相関を構成するのは何かを定義します)。

TIPS: データー診断と統計的な分析

- ストキャスティック・パラメーターの推定--統計的な分析とデーター診断では、ストキャスティック・パラメーターの推定上の表があり、ボラティリティ、ドリフト、平均回帰率とjump-diffusion率を履歴データーに基づいて推定します。これらのパラメーターの結果は使用された履歴データーだけに基づいており、適合された履歴的データーによって毎回パラメーターを変更しなければいけない点に気をつけてください。さらに、分析結果が全てのパラメーターを表示し、どのストキャスティック過程モデルが(例、ブラウン運動、平均回帰、Jump-Diffusion、あるいは混合過程)最も最適かどうかは問題となりません。どちらかと言うと予測する時系列変数によってユーザーが定義することです。分析は、ユーザーだけが知っている源からの最適な過程の定義を行いません(例、ブラウン運動糧は、株価のモデル化に最適ですが、分析は、分析された履歴的データーが株から、あるいは他の変数から定義されたものかどうかは定義できません。これは、ユーザーしか分かりせん)。最後に、ヒントとして、特定のパラメーターが通常範囲から外れている場合、この入力パラメーターが必要とする過程は、きっと適切な過程でない確率を持っているでしょう(例、平均回帰率が110%の場合、チャンスはあり、平均回帰

は適切な過程ではないなど)。

TIPS: 分配的な分析、グラフと確立表

- 分配的な分析--リスクシミュレーターで利用可能な42の確率分布のPDF、CDFとICDFを素早く計算し、これらの値の表を提供します。
- 分配的なグラフ表--同じ分布の異なったパラメーターを比較するために使用されます(例、例、[2, 2], [3, 5]と[3.5, 8]のアルファとベータを持ったワイブル分布の形状とPDF、CDFと ICDFの値と相互的なオーバーレイ)。
- オーバーレイグラフ--異なった分布を比較する為に使用され(理論的な入力仮定と実践的にシミュレーションされた予測結果)、表示的な比較のために相互的なオーバーレイをします。

TIPS: 効率的フロンティア

- 効率的フロンティアの変数--フロンティア変数のアクセスには、効率的フロンティアの変数の設定前にモデルの規制をまず設定します。

TIPS: 予測セル

- 公式なしの予測セル--公式、あるいは値を含まない予測結果のセルを設定することが可能(ただ警告メッセージを無視すること)だが、生成される予測結果のグラフが空白で表示されるかもしれないので注意してください。セルが頻繁に交信されたり、計算されるマクロスがある場合には、一般的に予測結果は空欄のセルに設定されます。

TIPS: 予測の結果グラフ

- タブ 対 スペースバー--いくつかの入力の記入後にキーボードのタブを打って予測の結果グラフを更新し、信頼値とパーセンタイルが得られます。また、スペースバーを打つと予測の結果グラフの様々な表が表示されます。
- ノーマル 対 グローバルビュー--表のインターフェースとグローバルインターフェースのどちらかに切り替えにこれらのビューをクリックし、予測の結果グラフの全ての要素が一度に表示されます。
- コピー--ノーマル、あるいはグローバルビューのどちらに切り替えてあるかによって、予測グラフ、あるいはグローバルビュー全体をコピーすることができます。

TIPS: 予測法

- セルのリンクアドレス--まずスプレッドシートでデーターを選択してから予測ツールを実行すると、選択されたデーターのセルアドレスは、自

動的にユーザーのインターフェースに入力され、そうでなければ、手動的にセルアドレスに入力するか、リンクアイコンを使用して重要なデータ配置に関連させなければいけません。

- 予測のRMSE--各モデルの確実性の直接的な比較のために複数の予測モデル上で予測のRMSEを一般的なエラー測定として使用してください。

TIPS: 予測法: ARIMA

- 予測周期--外因性データの行数は、時系列データの行を少なくとも希望する予測周期よりも上回らなければいけません(例、将来の5周期を予測をするのに100の時系列データのポイントを持っている場合、外因性変数上に少なくとも105、あるいはそれ以上のデータポイントが必要となります)。そうでなければ、制限なしで希望する周期と予測する外因性変数なしでただARIMAを実行します。

TIPS: 予測法: 基本的な計量経済学

- セミコロンの使った変数の区分--セミコロンを使用して従属変数を区分します。

TIPS: 予測法: ロジット、プロビットとトビット

- データの必要条件--ロジットとプロビットモデルの実行への従属変数は、2値でなければいけなく(0 と1)、トビットモデルは、2値と他の小数値を使用することができます。3つのモデルの全ての従属変数は、どの数値でも使用できます。

TIPS: 予測法: ストキャスティック過程

- デフォルト・サンプル入力--独自のモデルを進展するのに疑惑がある場合、デフォルト入力を初期ポイントとして使用します。
- パラメーター推定のための統計的な分析ツール--未加データからパラメーターの推定を行うことで、ストキャスティック過程モデル内の入力パラメーターを調節するツールとしてこれを使用します。
- ストキャスティック過程モデル--時々、ストキャスティック過程のユーザーインターフェースが長い間掛かるのは、入力に間違いがあり、モデルが正確に指定されないという好機です(例、平均回帰が110%の場合、好機は、平均回帰が適正な過程でないなど)。異なった入力、あるいは異なったモデルで再試行してみてください。

TIPS: 予測法: 傾向ライン

- 予測結果--レポートの下方にスクロールすると予測された値を見ることができます。

TIPS: 関数セル

- RSの関数--入力仮定の設定があり、Excelのワークシート内で使用できる予測統計の関数があります。これらの関数の使用には、まず RS 関数をインストールし(スタート、プログラム、リアル・オプションズ・バリユエーション、リスクシミュレーター、ツールと関数のインストールから)、Excel内にRS関数を設定する前にシミュレーションを実行。これらの関数の使用法の詳しい情報は、例証モデル24を参照。

TIPS: 練習とビデオ学習をはじめるにあたって

- 練習をはじめるにあたって--多数のステップごとの解説がふくまれた例証と結果の解釈のための練習がスタート、プログラム、リアル・オプションズ・バリユエーション、リスクシミュレーターのショートカット配置にあります。これらの練習は、ソフトウェアの使用に速く慣れるように考えられた物です。
- ビデオ学習をはじめるにあたって--全てのビデオは、www.realoptionsvaluation.com/download.html、あるいはwww.rovdownloads.com/download.htmlからダウンロードすることができます。

TIPS: ハードウェアのID

- 右クリックのHWIDコピー--ライセンスのインストールのユーザーインターフェースで、HWID上にある値を選択、あるいはダブルクリックし、右クリックでコピー、あるいは、Eメール HWID リンクをクリックしてHWIDを含んだメールを生成します。
- トラブルシューター--スタート、プログラム、リアル・オプションズ・バリユエーション、リスクシミュレーター・フォルダーとHWIDの取得ツールからトラブルシューターを起動。

TIPS: ラテン・ハイパーキューブ・サンプリング(LHS) 対 モンテカルロ・シミュレーション(MCS)

- 相関--入力仮定の間にペアワイズ相関を設定する時、リスクシミュレーターのオプションメニューからモンテカルロ・シミュレーションを使用することをお勧めします。ラテン・ハイパーキューブ・サンプリングは、シミュレーションのための相関されたコピュラ法と互換性を持っていません。
- LHS値--値の大きな数は、シミュレーションを遅くするのに対して、より一様なシミュレーション結果を提供します。
- ランダムネス--オプションメニューの全てのランダムシミュレーション法は検定されており、すべてが良いシミュレーターで、大きな試行回数

を実行するときにすべて同じレベルのランダムネスをアプローチします。

TIPS: オンラインの資源

- 書庫、ビデオ学習、モデル、ホワイトペーパー -- www.realoptionsvaluation.com/download.html、あるいは、www.rovdownloads.com/download.htmlから無料でダウンロードすることができます。

TIPS: 最適化

- 実行不可能な結果--最適化の実行が実行不可能な結果を提供した場合、規制を平等(=)から不平等(>= または <=)に変換して再実行することができます。これは、有効的なフロンティアの分析を実行しているときに適用できます。

TIPS: プロファイル

- 複数のプロファイル--単一モデルでの複数のプロファイルの作成と切り替え。これは、モデルの入力パラメーター、あるいは分布タイプを変換することによってシミュレーションでシナリオの実行を可能にし、結果への影響を見比べることができます。
- 必要とされるプロファイル--仮定、予測、あるいは決定変数は、アクティブなプロファイルがなければ作成することができませんが、利用可能なプロファイルがあれば、毎回新しいプロファイルを作成する必要はありません。実際に、仮定、あるいは予測を追加してからシミュレーションを実行したい場合、同じプロファイルを維持する必要があります。
- アクティブなプロファイル--Excelを起動すると、最後に保存されたプロファイルが自動的に呼び出されます。
- 複数のExcelファイル--Excelで開かれた複数のモデルの切り替えを行っている時、アクティブ・プロファイルは、その時に表示されているExcelモデルとなります。
- クロス・ワークブック・プロファイル--複数のExcelファイルを開いており、その中の1つだけがアクティブ・プロファイルの場合、誤って他のExcelファイルに切り替え、仮定や予測を入力して実行しても無効になります。これは、アクティブ・プロファイルでしか作業ができないからです。
- プロファイルの削除--既存するプロファイルの複製と既存するプロファイルの削除ができますが、Excelのファイルには、少なくとも1つのファイルがなければいけないことに注意してください。
- プロファイルの配置--作成したプロファイル(仮定、予測、決定変数、目

的、規制などを含んでいる)を暗号化された隠れたワークシートとして保存できます。従って、Excelワークブックファイルを保存すると自動的に暗号化されたプロファイルとして保存されます。

TIPS: 右クリックショートカットと他のショートカットキー

- 右クリック--Excelのどのセルでも右クリックからリスクシミュレーターのショートカットメニューを開くことができます。

TIPS: 保存

- Excelファイルの保存--プロファイルの設定、仮定、予測、決定変数とExcelのモデルの保存(リスクシミュレーターのレポート、グラフと抽出されたデータも含む)。
- グラフの設定の保存--予測の結果グラフの設定が保存できるように、同じ設定を回復し、将来の予測結果のグラフに適用することができます(予測グラフの保存と開くのアイコンを使用)。
- Excelでのシミュレーションデータの保存と抽出--シミュレーションされた実行の仮定と予測の抽出を行い、Excelファイルは、後ほどの回復の必要、あるいは可能性があるため、別に保存する必要があります。
- リスクシミュレーターでのシミュレーションされたデータとグラフの保存--リスクシミュレーター、データの抽出と*.RiskSim ファイルへの保存の使用によって、シミュレーションを再実行する必要がなく、同じデータが含まれたダイナミックで、生々しい予測のグラフを将来呼び出すことができます。
- レポートの保存と生成--シミュレーションのレポートと他の分析的なレポートは、ワークブックに区分されたワークシートとして抽出され、Excelファイル全体が後ほどの呼び出しに備えて保存する必要があります。

TIPS: サンプリングとシミュレーション法

- 乱数の生成--サポートされた6つの乱数の生成が(詳細にはユーザーマニュアルを参照)あり、一般的にROV リスクシミュレーターのデフォルト法と高度な減算乱数シャッフル法の2法の使用がお勧めのアプローチです。モデル、あるいは分析が他の技法を特別に必要としない限り他の技法を適用しないで下さい。また、上記で挙げた2つの技法の結果の検定を行う事もお勧めします。

TIPS: ソフトウェアの開発キット(SDK)とDLL ライブラリー

- SDK、DLLとOEM--リスクシミュレーターの全ての分析はこのソフトウェア以外で呼び出すことができ、他のソフトウェアに統合することができます。開発キットソフトウェアの使用の詳細には、contact

admin@realoptionsvaluation.com へご連絡ください。ダイナミック・リンク・ライブラリー(DLL)の分析ファイルにアクセスできます。

TIPS: Excel でリスクシミュレーターをはじめるにあたって

- ROV のトラブルシューター--このトラブルシューターを実行することでコンピュータのHWIDが取得でき、ライセンスのためのコンピュータの設定、予備必要条件を見ることができ、誤ってリスクシミュレーターの利用を不可能にしてしまった場合に、再利用を可能にすることができます。
- Excelの起動時のリスクシミュレーターの開始--Excelの起動時にリスクシミュレーターが自動的に開始させることができます。また、手動的にスタート、プログラム、リアル・オプションズ・バリュエーション、リスクシミュレーター・ショートカット配置からアクセスすることもできます。これは、リスクシミュレーターのオプションメニューから設定することができます。

TIPS: 超高速シミュレーション

- モデルの開発--モデルで超高速シミュレーションを実行したい場合、モデルの構成時にいくつかの超高速シミュレーションを試験し、モデルの出来上がりに超高速シミュレーションの実行に問題が出ないことを保証していきます。超高速シミュレーションの検定を行うのにモデルが完全に出来上がるのを避けてください。後から壊れたリンクや互換性のない関数を見出す作業の法が困難です。
- 普通のスピード--モデルに疑惑を持っている場合、普通のスピードのシミュレーションが使用できます。

TIPS: 竜巻分析

- 竜巻分析--竜巻分析を1度だけ実行してはいけません。これは、モデルの診断ツールとしての意味で、同じモデル上で何回でも実行するのが理想だということです。例えば、大きなモデルですべてをデフォルト設定にして1回目の竜巻を実行し、すべての先例を表示します(全ての変数の表示を選択)。このシングル分析からは、大きくて、長いレポート(見にくい)の竜巻グラフが表示されることでしょう。しかし、これが基準となる初期ポイントとなり、先例の重要な成功要素数の考慮を定義し(例、竜巻グラフが表示する最初の5つの変数は、結果に高いインパクトを持っていることを示すのに対して、残りの200の変数は、小さい影響を示すことを表す)、どのケースにしても、2番目の竜巻分析は、より少ない変数で実行することができます(例、上位5が重要な変数を示すなら、上位10の変

数の表示を選択することで、竜巻グラフは、重要な要素とそうでもない要素の比較を表示し、対照が含まれた良いレポートを作成することができます。従って、竜巻分析には、重要な要素とそうでもない要素の極値を含み、結果への影響を表示することを勧めます。).

- デフォルト値--デフォルト検定のポイントは、 $\pm 10\%$ から非線形性のため大きな検定値にまで増加することができます(スパイダーグラフは、非線形ラインを表示し、先例の影響が非線形である場合、竜巻グラフは片側への歪みを表示します)。
- ゼロ値と整数--ゼロ、あるいは整数値の入力は、実行前に竜巻分析のみで選択をはずさなければいけません。そうでなければ、摂動の%がモデルを無効にします(例、モデルが索引表を使用し、1月 = 1、2月 = 2、3月 = 3...とした場合、1の値を $\pm 10\%$ で摂動するため、0.9 と 1.1 という結果が表示され、モデルの意味がなくなってしまう)。
- グラフのオプション--モデルへの最的なグラフタイプを見出すためにさまざまなグラフオプションを試してください。

TIPS: トラブルシューター

- ROVのトラブルシューター--このトラブルシューターを実行することでコンピューターのHWIDが取得でき、ライセンスのためのコンピューターの設定、予備必要条件を見ることができ、誤ってリスクシミュレーターの利用を不可能にしてしまった場合に、再利用を可能にすることができます。