

리스크 시뮬레이터 2012

사용자 매뉴얼



본서는 Risk Simulator 2012 의 사용에 대한 이해를 돕기 위하여 구성되어 있습니다. 본서에 기재된 화면은 실제 화면과 다를 수 있으며, 본서 및 본서 내에 기재되어 있는 소프트웨어는, Real Options Valuation 주식회사의 한국 총판인 (주) 마이크로폴리스가 Real Options Valuation 주식회사로부터 사용 허가 계약을 기반으로 제공되고, 사용 승낙 계약 조항에 따라서 사용 혹은 복사하는 것이 가능합니다. 본서에 기재된 내용은 정보의 제공만을 목적으로 하고 있어, 예고없이 변경 될 수 있습니다. (주) 마이크로폴리스와 Real Options Valuation 주식회사는, 본서의 내용에 대하여 어떠한 책임도 지지않습니다.

사용 승낙 계약 조항에서 허가되어 있는 경우를 제외하고, 본서의 일부 혹은 전부를 (주)마이크로폴리스와 Real Options Valuation 주식회사의 사전의 서면에 의하여 허락없이, 전자적, 기계적, 녹음을 포함한 어떠한 수단 및 형식에 의하여도, 복제, 검색 시스템에의 보존, 또는 전송하는 것은 금지되어 있습니다.

본서는 Real Options Valuation 주식회사의 창업자 겸 CEO인 조나단 문에 의한 저작의 사용 승락 계약에 기초합니다.

저작권 : 조나단 문

저작, 디자인 및 인쇄 : 미국

© 2005-2012 Dr. Johnathan Mun. All rights reserved. www.realoptionsvaluation.com

Microsoft® 는, 마이크로 소프트 코퍼레이션의 미국 및 다른 국가에서의 상표 등록입니다. 기타 모든 상표는, 해당하는 각 사가 소유하고 있습니다.

목차

1.		7
	?	7
		8
		9
2011/2012	Risk Simulator	12
2.		18
	?	18
		19
	소프트 웨어의 고도의 전략	19
	몬테카를로 시뮬레이션의 실행에 대하여	19
	1.	20
	2. 가	22
	3.	24
	4. 시뮬레이션 실행.....	25
	5. 예측 결과의 해석	25
	상관의 기초.....	32
	리스크 시뮬레이션에서 상관을 적용하는 것에 대하여.....	33
	몬테카를로 · 시뮬레이션에서의 상관의 영향.....	34
	분포의 중심을 측정하는 것에 있어서 —제 1 차 모멘트.....	38
	분포의 파급의 측정 —제 2 차 모멘트	38
	분포의 왜도를 측정하는 것에 있어서—제 3 차 모멘트.....	40
	분포에서의 파국적인 Tail Event 를 측정하는 것에 있어서 —제 4 차 모멘트.....	41
3.		42
	가	42
		44
		45
		49
		53

.....	55
Box-Jenkins ARIMA	58
Auto ARIMA (Box-Jenkins ARIMA).....	63
.....	64
J-S	65
GARCH	66
(Markov Chains)	69
(MLE: Logit, Tobit, Probit).....	70
Spline (Cubic Spline Interpolation and Extrapolation/입방 스플라인).....	72
4.	74
.....	74
.....	76
.....	82
5.	87
.....	87
.....	94
:	97
Boot Strap	102
가	105
.....	107
.....	108
.....	110
.....	119
.....	123
Risk Simulator 2011/2012	128
,	128
Deseasonalize Detrend	129
.....	131
.....	132

<i>Trendline</i>	134
	135
	136
:	137
<i>ROV BizStats</i>	142
	147
<i>Goal Seek (Optimizer Goal Seek)</i>	152
(Single Variable Optimizer /)....	153
(Genetic Algorithm Optimization).....	154
ROV	157
<i>Decision Tree</i>	157
<i>Simulation Modeling</i>	159
<i>Bayes Analysis</i>	160
<i>Expected Value of Perfect Information, MINIMAX and MAXIMIN Analysis, Risk Profiles, and Value of Imperfect Information</i> 가 , , , 가	160
<i>Sensitivity</i>	161
<i>Scenario Tables</i>	161
<i>Utility Function Generation</i>	161
	170
: 가 (가)	170
: /	170
:	171
:	171
:	172
:	172
:	172
:	173
:	173
: : ARIMA	173

:	: Basic Econometrics.....	173
:	: Logit, Probit, and Tobit	174
:	: Stochastic Processes	174
:	: Trendlines	174
:		174
:		175
:	ID	175
:	(LHS) (MCS).....	175
:		175
:		176
:		176
:	가	177
:		177
:		177
:	(SDK) DLL	178
:	Risk Simulator	178
:		178
:	Tornado Analysis	178
:	Troubleshooter	179

1.

?

리스크 시뮬레이터 (RiskSim)의 1.1 버전은 몬테카를로 · 시뮬레이션을 사용한, 예측 및 최적화를 위한 소프트웨어입니다. 이 소프트웨어는 Microsoft .NET C#으로 구성되어 있어, Add-on 된 Excel 과 함께 사용합니다. 이 소프트웨어는 호환성이 있으므로, Real Options Valuation, Inc.의 소프트웨어 Real Options Super Lattice Solver (SLS)및, Employee Stock Options Valuation Toolkit (ESOV)과도 사용이 가능합니다.

이 매뉴얼은 사용자에게 도움이 되는 충분한 정보를 기재하기 위하여 노력하였으나, 소프트웨어의 개발자 (예를 들어, 조나단 문 저작의 리얼 옵션 분석, 제 2 판, Wiley Finance 2005 ; 리스크 모델화: 몬테카를로 · 시뮬레이션의 적용, 리얼 옵션 분석, 예측, 그리고, 최적화, 제 2 판, Wiley 2006 ; 그리고, 종업원 주식 구입 선택권의 평가 (2004 FAS 123R), Wiley Finance 2004 의 트레이닝 DVD, 라이브 트레이닝 및 본서등의 대용이 되지 않는 것의 양해를 바랍니다. 이 상품에 대하여 상세 정보는 홈페이지를 참고하십시오.

www.micropolis.co.kr (South Korea) www.realoptionsvaluation.com (USA) .

리스크 · 시뮬레이터의 소프트웨어에는 다음의 모듈이 있습니다.

- 몬테카를로 · 시뮬레이션 (Parametric 및 Nonparametric 시뮬레이션으로 42개의 분포 확률과 동시에 여러 가지 시뮬레이션 프로파일의 실행. 혹은 절단, 상관 시뮬레이션, 커스터마이즈 시뮬레이션, 정확함과 에러를 컨트롤하는 시뮬레이션, 기타, 알고리즘을 사용한 시뮬레이션 등의 실행) .
- 예측 방법 (Box-Jenkins ARIMA, 중회귀 분석, 비선형 외압법, 확률 과정, 그리고 시계열 분석의 실행) .
- 불확정 상황에서의 최적화 (시뮬레이션의 실행, 무실행과는 무관계로 포트폴리오 및 프로젝트의 최적화를 위하여 이산 정수 및 연속 변수를 사용하는 최적화를 실행) .
- 모델화와 분석 툴 (Boot Strap · Simulation, 가설 실험, 분포 적합등과 같이 토네이드, 스파이더 및 감도 분석의 실행) .

Real Options SLS 의 소프트웨어는 단일, 혹은 복합적 옵션을 컴퓨터에 측정하는 이외에, 사용자의 니즈에 대응하여 조절 가능한 옵션 모델을 사용하는 것이 가능합니다. 이 소프트웨어에는, 다음과 같은 모듈이 있습니다.

- 단일 자산 SLS (커스터마이즈 옵션 해결과 같이, 포기, 선택자, 단축, 연기 및 확장의 옵션등의 해결) .
- 복합 자산과 복합 변화 SLS (지속적인 복합 변수의 해결, 근본적인 자산 및 다단계의 해결 옵션, 포기, 선택자 단축, 연기 및 확장등과 지속적인 다단계의 결합, 그리고 변환의 옵션. 즉 커스터마이즈 옵션의 해결도 사용할 수 있습니다) .

- 다항의SLS (삼항식 (평균-회귀) 옵션, 사항식 (jump-diffusion) 옵션 그리고, 5항식 (레인보우) 옵션의 해결)。
- Excel Add-on기능 (모든 전례의 옵션의 해결의 다른 패형식 모델이 더해진 해결 및, Excel에 기초한 환경내에서의 커스터마이즈 옵션의 해결)。

이 소프트웨어를 인스톨하기 위하여는, 화면상의 가이드를 따라 주십시오. 즉 최소한의 요구 시스템이 준비되지 않으면 인스톨이 되지 않는 것을 알아 주시기 바랍니다。

- Pentium IV 프로세서 및 이상
- Windows XP 또는, Vista, Windows 7
- Microsoft Excel XP, 2003, 2007, 2010 및 이상
- Microsoft .NET Framework 2.0/3.0.
- 500MB 의 공간 확보
- 2GB RAM 를 추천
- 소프트 웨어·인스톨의 관리자 권한

새로운 컴퓨터의 대부분은 Microsoft .NET Framework 이 이미 Pre-Install 되어 있습니다. 리스크·시뮬레이터의 인스톨을 실행하기 위하여 .NET Framework 의 인스톨이 필요하다는 에러 메시지가 표시되는 경우에는, 인스톨을 중단하고, 인스톨 CD 에서 적절한 .NET Framework 을 인스톨하여 주십시오 (언어 선택을 잊지 마십시오) 。 .NET Framework 의 인스톨이 끝나고, 컴퓨터를 재부팅한 후에, 리스크 시뮬레이터 S/W 의 인스톨을 다시 실행하여 주십시오. .NET Framework 2.0 / 3.0 의 어떤 버전에서도 컴퓨터에 이미 인스톨되어 있지 않은 리스크 시뮬레이터가 실행되지 않는 것을 이해 바랍니다。

30 일간의 Trial 무료 라이선스의 파일이 소프트웨어에 부착되어 있습니다。

이후, 정규판을 구입하는 경우에는, Real Options Valuation 주식회사의 한국 총판인 (주) 마이크로폴리스의 메일 rov@micropolis.co.kr 사이트, www.micropolis.co.kr, 혹은 02) 563-8884 의 전화번호로 연락을 주십시오. www.realoptionsvaluation.com (USA) 당사의 Web 사이트에서는, 최신의 소프트웨어의 다운로드가 가능합니다. 또한, FAQ 링크에서는 라이선스에 관한 최신 업 데이트 정보 및, 인스톨 시에 일어나는 문제에 대하여 해결 방법을 제공할 수 있습니다。

소프트웨어의 정규판을 구입하는 경우에는, 고객의 라이선스 파일을 작성하기 위하여, 하드웨어의 ID 를 본사에 메일로 보내주시기 바랍니다. 다음의 항목에 따라 주시기 바랍니다.

Windows XP 시스템에서, Excel XP/2003/2007/2010 를 사용하는 경우:

- 머리말, Excel 에서 **리스크 시뮬레이터 (Risk Simulator) | 라이선스 (License)** 를 클릭하고, 화면에 표시 된 1 1 -20 자리의 코드와 영수학의 하드웨어 ID 를 카피하여 , rov@micropolis.co.kr 앞으로 메일을 보내주시기 바랍니다. (하드웨어 ID 를 선택하여 우클릭하여 복사하거나, e-mail 하드웨어 ID 링크를 클릭하여 주십시오.) 나중에, 본사에서 구입하신 소프트웨어의 영구 라이선스 ID 를 고객에게 메일로 송부합니다. 이 영구 라이선스 ID 의 파일은 고객의 컴퓨터 하드 드라이브에 보존하여 주십시오. 그리고 Excel 을 실행한 후에, **리스크 시뮬레이터 (Risk Simulator) | 라이선스 (License)** 를 클릭하여, **Install License (라이선스의 인스톨)** 을 클릭하여, 영구 라이선스 ID 의 파일을 선택하면 완료됩니다. 이것으로, 고객은 걱정없이 정규판의 소프트웨어를 실행하는 것이 가능합니다. 이 프로세스에 걸리는 시간은 1 분 이내입니다.

Windows Vista/Windows 7 시스템에서, Excel XP/2003/2007/2010 을 사용하는 경우:

- 처음에, Excel 2007/2010 을 Windows Vista/Windows 7 내에서 실행하여 주십시오. 리스크·시뮬레이터의 메뉴바에서 라이선스 아이콘을 클릭하거나, **리스크 시뮬레이터 (Risk Simulator) | 라이선스 (License)** 를 클릭하여, 화면에 표시되는 1 1 -20 자리의 코드와 영·수 문자의 하드웨어 ID 를 카피하여, rov@micropolis.co.kr 앞으로 메일을 하여 주십시오. (하드웨어 ID 를 선택하여 우클릭하고 카피하거나, e-mail 하드웨어 ID 링크를 클릭하여 주십시오.) 이후에, 본사에서 구입하신 소프트웨어의 영구 라이선스 ID 가 고객에게 메일로 송부됩니다. 이 영구 라이선스 ID 의 파일은 고객의 컴퓨터의 하드 드라이브에 보존하여 주십시오. 그 후에, 다음 스텝을 행하여 주십시오. Excel 을 실행하여, **리스크 시뮬레이터 (Risk Simulator) | 라이선스 (License)** 또는, 라이선스 아이콘을 클릭하여, **Install License (라이선스의 인스톨)** 을 클릭하고, 영구 라이선스 ID 의 파일을 선택하여 주십시오. 그리고, Excel 을 재실행하면 완료 됩니다. 이것으로, 고객은 걱정 없이 정규 소프트웨어를 실행하는 것이 가능합니다. 이 프로세스에 걸리는 시간은 1 분 이내입니다.

인스톨 완료 후, Microsoft Excel 을 재실행하여 주십시오. 인스톨 프로세스가 완료한 경우에는 Excel XP/2003 의 메뉴바에 추가의 리스크·시뮬레이터가 표시됩니다.

또한, Excel 2007/2010 의 경우는 Add In 그룹내에 표시되어, 새로운 아이콘 바가 Excel 에 표시됩니다. (Figure1.1 를 참조하여 주십시오.)

또한, 스프래쉬 스크린이 표시되어, 소프트웨어가 Excel 내에서 기능하고 있는 증거입니다. (Figure1.2) Figure1.3 은 리스크 시뮬레이터의 툴 바입니다. 전열의 아이템이 Excel 에 표시되어 있으면, 소프트웨어를 바로 실행 할 수 있습니다. 다음 장에서는 소프트웨어의 사용법이 자세히 기재되어 있습니다.

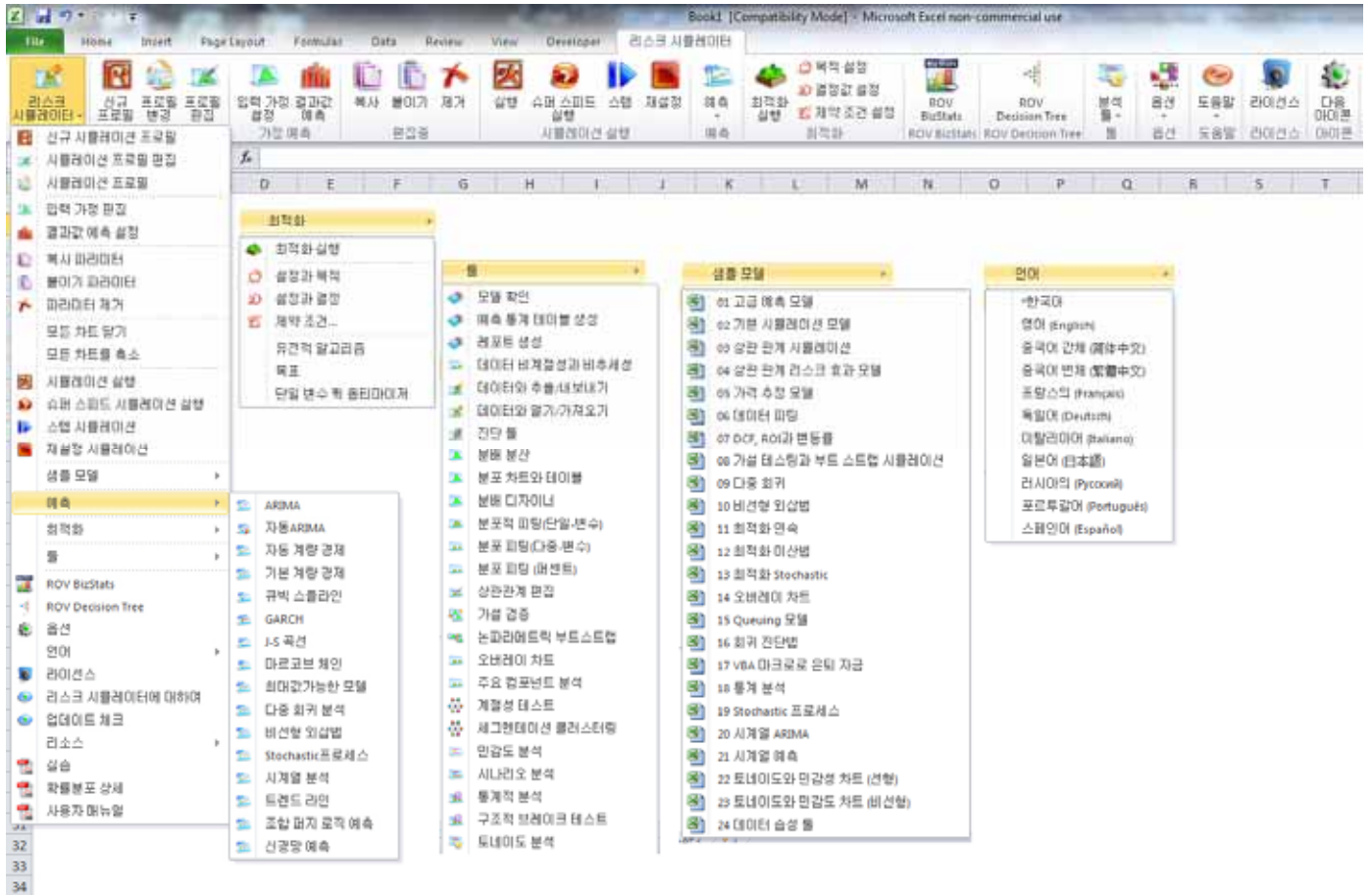


Figure 1.1 – Excel 2007/2010 의 리스크 시뮬레이션 메뉴와 아이콘 바



Figure 1.2 – 리스크·시뮬레이터 스프래쉬 스크린

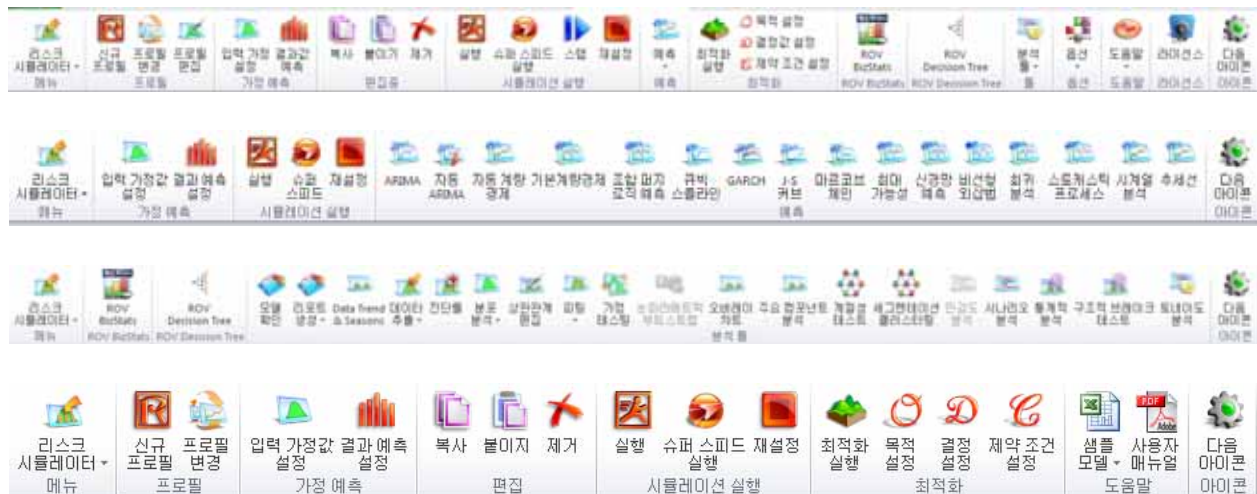


Figure 1.3 – Excel 2007/2010 의 리스크 시뮬레이터 툴바의 아이콘

Risk Simulator

1. 11 — 1 1 1 1 1 2 2
2 1 2 1 2 2 2

-
-
-
-
- ()
-
- MINIMAX
- MAXIMIN
-

5. —Risk Simulator 24

300

9. —Window 7, Vista, XP ; 2010, 2007, 2003
; virtual machine MAC OS

© 2005-2012 Real Options Valuation, Inc.

11. — Risk Simulator
가
12. / —가 , , /
13. — (가 , , ,)
14. 2007/2010 — 가 (1280 x 760).
15. - — Risk Simulator
16. ROV —Real Options SLS, Modeling Toolkit, Basel Toolkit, ROV Compiler, ROV Extractor and Evaluator, ROV Modeler, ROV Valuator, ROV Optimizer, ROV Dashboard, ESO Valuation Toolkit ROV
17. RS —가 RS
18. : , , ID
19. : (5.2).).
20. , , — , , ,

Simulation Module

21. 6 가 —ROV Advanced Subtractive Generator, Subtractive Random Shuffle Generator, Long Period Shuffle Generator, Portable Random Shuffle Generator, Quick IEEE Hex Generator, Basic Minimal Portable Generator
22. 2 가 —Monte Carlo Latin Hypercube.
23. 3 가 Correlation Copulas— Normal Copula, T Copula, and Quasi-Normal Copula .
24. 42 가 —Arcsine, Bernoulli, Beta, Beta 3, Beta 4, Binomial, Cauchy, Chi-Square, Cosine, Custom, Discrete Uniform, Double Log, Erlang, Exponential, Exponential 2, F Distribution, Gamma, Geometric, Gumbel Max, Gumbel Min, Hypergeometric, Laplace, Logistic, Lognormal (Arithmetic) and Lognormal (Log),

Lognormal 3 (Arithmetic) and Lognormal 3 (Log), Negative Binomial, Normal, Parabolic, Pareto, Pascal, Pearson V, Pearson VI, PERT, Poisson, Power, Power 3, Rayleigh, T and T2, Triangular, Uniform, Weibull, Weibull 3.

25. — .
26. —
27. —
28. — 가
29. —
30. —
31. —100,000
32. ARIMA— ARIMA (P,D,Q).
33. ARIMA—가 가 ARIMA
34. — 가
(linear, nonlinear, interacting, lag, leads, rate, difference).
35. — , / ,
36. —
37. GARCH— volatility projection:
GARCH, GARCH-M, TGARCH, TGARCH-M, EGARCH, EGARCH-T, GJR-GARCH, GJR-TGARCH
38. J- — J
39. —Logit, Probit, Tobit
40. —
41. — (forward, backward, correlation, forward-backward).
42. —
43. S — S
44. — , , 8 가
45. Trendlines— linear, nonlinear polynomial, power, logarithmic, exponential, moving average .
46. (, , ,)
- 47.

48. —
49. — Hessian matrices, LaGrange function
50. —
51. —
52. — , , , ,
53. —
54. —
55. — 가
56. — 가
57. — ,
58. —
59. **Check Model**— 가
60. Correlation Editor—
61. Create Report— 가
62. Create Statistics Report—
63. Data Diagnostics— heteroskedasticity, micronumerosity, outliers, nonlinearity, autocorrelation, normality, sphericity, nonstationarity, multicollinearity, correlations
64. Data Extraction and Export— , , Risk Sim
65. Data Open and Import—
66. **Deseasonalization and Detrending**— deseasonalize detrend
67. Distributional Analysis—42 가 PDF, CDF, ICDF
68. Distributional Designer—

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.

몬테카를로 · 시뮬레이션이란, 모나코의 유명한 도박 도시의 이름에서 나온, 매우 유력한 방법론으로 알려져 있습니다. 실무가로서, 시뮬레이션은 현실에서 일어난 복잡하고 곤란한 문제를 해결로 이끌어 줍니다. 몬테카를로는 수 천, 혹은 수백만의 확률 샘플 경로에 기초하여 인공적인 미래의 예상도를 그려 내어 줍니다. 회사에서의 분석은, 대학원의 고등수학을 수강하는 것만으로 이론적이지도 실천적이지도 않습니다. 우수한 애널리스트는, 자유롭게 모든 수법을 사용하여, 가능한한 무엇보다도 용이하고 또 실천적인 방법으로 같은 결과를 얻는 사람들을 말합니다. 어떠한 케이스에서도, 그 모델이 정확한 경우, 몬테카를로 시뮬레이션은 유의한 답변에 무엇보다도 수학적 방법을 부여합니다. 그러면 무엇보다 자세한 몬테카를로 시뮬레이션에 대하여, 그리고 어떠한 상황에서 사용할 것인지 봅시다.

?

몬테카를로 · 시뮬레이션이란, 예측, 추정, 리스크의 분석 등에 도움이 되는 난수를 이용하는 무엇보다도 단순한 시뮬레이션법입니다. 시뮬레이션은, 불확정 변수에 대하여 사용자가 사전에 정한 확률 분포에서 값을 반복적으로 채취하거나, 모델의 값을 사용하여 모델의 몇 개인가를 계산합니다. 전체적으로, 각 시나리오가 각각의 예측값을 사용하여 얻을 수 있도록 모델로서, 이것들의 시나리오는 관련될 결과를 만들어 내고 있습니다.

예측이란, 사용자가 모델의 중요한 산출로서 정의한(대부분, 공식 및 관수를 사용하여)사상으로, 총 수입, 순이익, 지출 총액 등을 포함합니다.

골프 볼을 커다란 바구니에서 교대로 몇 번이고 꺼내는 상황을 생각 해 주시기 바랍니다. 먼저, 바구니의 형태와 사이즈에 의하여 분포의 입력 가정 *input assumption* (평균 100 및 표준편차 10의 정규 분포와 *v s* 동일 분포 혹은, 삼각 분포) 가 결정됩니다. 즉, 바구니에 따라서는 보다 바닥이 깊은 것보다 대칭적인 것이 있어, 그 가정에 의하여 어떠한 정도의 빈도로 다른 것보다 확실히 볼을 꺼낼 수 있는 것이 가능하게 됩니다.

반복적으로 꺼내지는 볼의 수는, 시뮬레이션한 회수 (*trials*) 에 의하여 다르게 됩니다. 복수의 관련된 가정을 가진 규모의 커다란 모델의 경우는, 커다란 바구니중에 작은 바구니가 몇 개 들어 있어, 이것의 작은 바구니는 각각에 골프볼의 세트가 있어 여기 저기 날고 있는 상황을 생각하여 주십시오. 예를 들어 변수의 사이에 상관 관계가 있다고 하면 그것의 바구니를 손에 들어 맞추는 것에 의하여, 다른 바구니의 골프 볼과는 다르게, 상호 관련된 바구니의 골프 볼은 종렬로 서서 날아가는 것이 됩니다. 이 상호 작용에서 픽업 된 볼 하나 하나는 시뮬레이션의 예측 결과 (*forecast output*) 로서 기록되고, 표에 기록 됩니다.

소프트웨어의 고도의 전략

리스크·시뮬레이션 소프트웨어에서는 몬테카를로·시뮬레이션, 예측, 최적화등의 여러 가지 어플리케이션이 포함되어 있습니다.

- ② 이 시뮬레이션 어플리케이션은 Excel 에 기초한 모델, 예측의 시뮬레이션 (분포 결과) 、적합 분포의 표시 (무엇보다 적합한 통계적 분포를 자동적으로 보여 주고 있습니다) 、상관의 컴퓨터화 (변수의 사이에서의 관계 유지) 、감도의 구분 (토네이도, 감도표 작성) 등을 커스터마이징 및 Non Parameter 시뮬레이션과 같이 (분포의 상세와 그것의 파라미터를 제외한 이력 데이터를 사용한 시뮬레이션) 실행합니다.
- ② 예측 어플리케이션에서는, 계량 경제 예측 Box-Jenkins ARIMA, 자동 시계열 예측 (계절성과 트렌드와 같이) 다변량 회귀 (선형, 비선형 회귀) 、비선형 외압법 (곡선 적합) 그래서, 확률 과정 (랜덤 워크, 평균 회귀, jump-diffusion 과정) 을 만들어 내는데 사용됩니다.
- ② 최적화 어플리케이션은, 목적치를 최대, 최소로 하는 제약 하에, 중이산적 정수, 연속, 혼합 결정 변수의 최적화 및 、몬테카를로·시뮬레이션과 함께 정지·동적·확률 최적화 등의 실행, 그리고 선형, 비선형의 최적화에도 사용할 수 있습니다. Real Options Super Lattice Solver 는 Stand Alone S/W 로 리스크 시뮬레이터를 보완하여, 단순 혹은 복잡한 리얼 옵션의 문제를 해결하여 줍니다.

몬테카를로 시뮬레이션의 실행에 대하여

통상, Excel 모델에서 시뮬레이션을 실행할 때에, 다음의 순서를 행하여 주십시오.

1. 신규 시뮬레이션 프로필을 작성하거나, 미리 보존하고 있는 프로필을 열어 주십시오.
2. 입력 과정의 설정을 적용한 셀에서 행하여 주십시오.
3. 예측 결과의 설정을 적용한 셀에서 행하여 주십시오.
4. 시뮬레이션을 실행하여 주십시오.
5. 결과를 해석하여 주십시오.

예 : 시뮬레이션의 실행 사례를 시험하기 위하여 *Basic Simulation Model* 이라는 파일을 열어 주십시오. 이 파일은 스타트 메뉴에서도 여는 것이 가능합니다. 다음의 순서에서 파일을 클릭하여 주십시오. *Start / Real Options Valuation / Risk Simulator / Examples* 또는 *Risk Simulator / Example Models* 에 직접 들어가 주십시오.

1.

신규 시뮬레이션을 실행하는 것은 먼저, 신규 시뮬레이션 프로파일의 작성이 필요합니다. 시뮬레이션 프로파일은 시뮬레이션을 실행하기 위하여 필요한 (가정, 예측, 실행의 환경 설정 등) 순서의 모든 설정이 표시됩니다. 프로파일을 작성하는 복수의 시뮬레이션 시나리오를 설정 할 수 있습니다. 이것은 복수의 사용자가 같은 모델 (개인적인 설정을 지정하고) 을 사용할 수 있다는 것입니다. 또는, 한 사람의 사용자가 같은 모델에서 몇 개의 시나리오 예측을 다른 분포 가정등에서 검색하는 것도 할 수 있습니다.

- Excel 을 열고, 신규, 또는 이미 보존하고 있는 모델을 열어 주십시오. (예증의 경우는, 모범 모델의 기초 시뮬레이션을 열어 주십시오)。
- *Risk Simulator* 를 클릭하고, *New Simulation Profile* 을 선택하여 주십시오.
- 시뮬레이션에 적절한 타이틀명을 기입하여 주십시오. (Figure 2.1 참조)

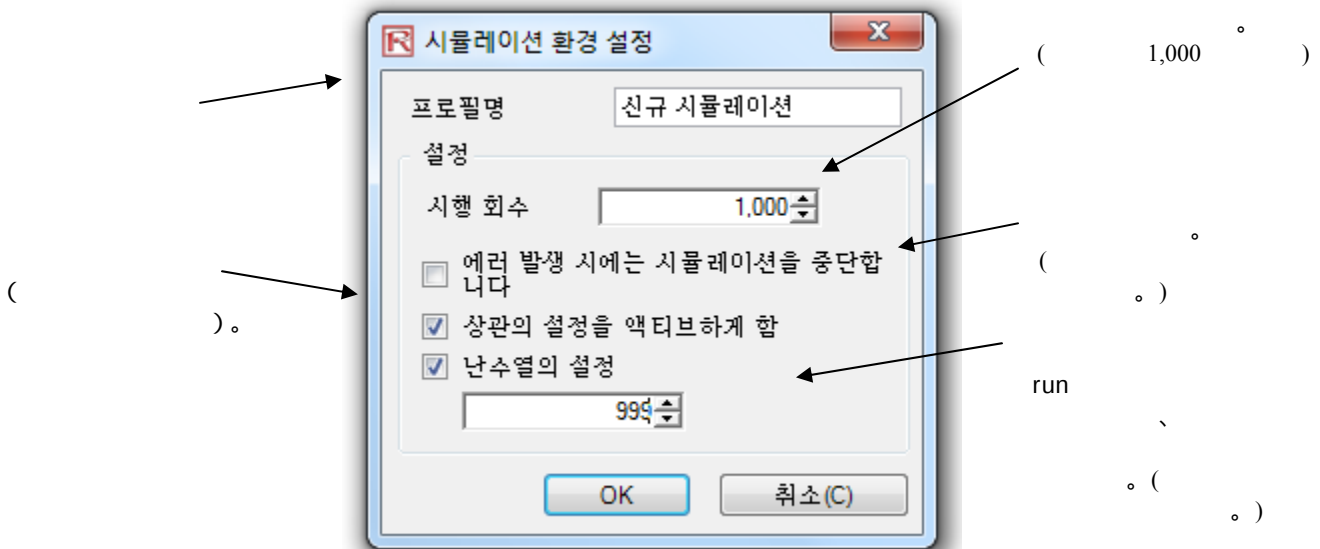


Figure 2.1 - 신규 시뮬레이션 프로파일

- ① **타이틀** : 시뮬레이션 타이틀을 지정하는 것으로, 하나의 Excel 모델에 몇 개의 시뮬레이션 프로파일을 작성하는 것이 가능합니다. 즉, 하나의 모델에서 몇 개의 시뮬레이션 시나리오를 보존하는 것이, 전 회에 사용한 가정 및 신규 시나리오의 편집을 행할 필요가 없이 가능합니다. 언제라도 프로파일 명칭을 리스크 시뮬레이터 (*Risk Simulator*) / *프로파일의 편집 (Edit Profile)* 을 선택한 후에 편집하는 것이 가능합니다.
- ② **시행수** : 시뮬레이션이 시행된 회수의 지정을 행할 때에 표시됩니다. 즉, 1,000 시행은 입력 가정에 기초한 1,000 의 다른 출력 결과가 생성된다는 것입니다. 시행 회수는 사용자의 희망으로 자유롭게 편집하는 것이 가능합니다만, 입력 시에 제대로 된 정수가 아니면 않되는 것을 잊지

마십시오. 디폴트에서 이미 1,000 회로 표시됩니다. 또는, 정도 와 에러 컨트롤을 사용하는 것 하면 자동적으로 시뮬레이션에 맞는 필요한 시행의 횟수가 결정됩니다. (상세는 정도와 에러 컨트롤의 셀렉션을 보시기 바랍니다) 。

- ② 시뮬레이션 에러 시에 일시 정지 : 선택을 하면 Excel 모델에서 에러가 발생하는 때에 항상 시뮬레이션을 정지합니다. 예를 들어, 모델에 에러가 생길 경우 (시뮬레이션 시행으로 생성된 입력 값은, 어느쪽인가의 분산 시트 셀에서 에러로서 제로로 나누기가 되어 있어야 합니다) 는 시뮬레이션이 정지됩니다. 이것은 Excel 모델상에서 컴퓨터 에러가 없는가를 확인하기 위하여 필요한 항목입니다. 따라서, 모델의 구성에 문제가 없으면 미리 알고 있을 때에 이 툴을 사용할 필요는 없습니다.
- ③ 상관의 체크 : 체크하면 짝이 되어 있는 input 가정의 상관이 계산됩니다. 그렇지 않으면, 상관은 모두 제로로 설정되고, 시뮬레이션의 실행 시, 상관과 입력 과정 사이의 값은 계산되지 않습니다. 예를 들어, 정말 상관 관계가 존재한다고 하면, 상관과 입력 가정의 계산을 행하는 것으로 더욱 확실한 결과가 나옵니다, 그리고 마이너스의 상관이 있다고 하면, 낮은 신뢰도의 예측이 나타날 경향이 보여 집니다. 이 설정에서 상관의 박스를 체크한 후, 생성된 각 가정의 상관 계수의 설정이 됩니다 (상세는 상관에 대한 장을 보시기 바랍니다) 。
- ④ 난수열의 지정 : 난수열을 지정한 시뮬레이션은 매회 미묘하게 다른 결과를 보여 줍니다. 이것은 몬테카를로·시뮬레이션의 난수 생성열의 기능에 의하여, 모든 난수 생성열에서는, 이론상의 사실입니다. 그러나, 프레젠테이션을 할 때에는 이미 인쇄 된, 또는 표시된 결과와 같은 결과를 표시하지 않으면 않습니다. 그 때에는 이 항목을 체크하고, 초기 시트 수를 입력하여 주십시오. 시트 수는 플러스 정수를 사용하여 주십시오. 같은 초기 시트 수, 같은 시행 수, 그리고 같은 입력 가정을 기입하는 시뮬레이션은 이미 같은 난수열을 선택하고, 같은 결과가 나오는 것을 보증합니다.

시뮬레이션 프로필이 한번 생성된 후에도, 이 속성을 편집할 수 있는 것을 잊지 마십시오. 이것을 실행하기 위하여는 먼저, 열려 있는 프로필이 편집하고자 하는 프로필인가를 확인하십시오. 또는, [리스크 시뮬레이터 \(Risk Simulator\) / 시뮬레이션 프로필의 편집 \(Change Simulation Profile\)](#) 를 클릭하고, 편집하고자 하는 프로필을 선택하여, **OK** 를 클릭하여 주십시오. (Figure 2.2 에서는 복수의 프로필 중에서 편집하고자 하는 프로필의 선택 화면을 참조하고 있습니다) 。 덧붙여 프로필의 카피, 명칭 변경도 가능합니다.

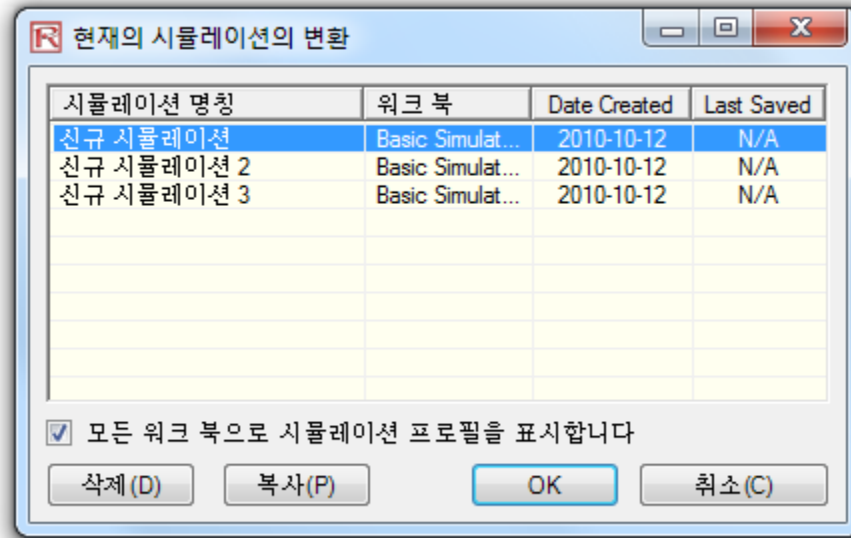


Figure 2.2 – 현재 액티브한 시뮬레이션의 변경

2. 가

다음의 단계는 모델의 입력 가정의 설정입니다. 가정은 셀에 방정식 및 관수를 제외한값을 기입하여 주십시오. 예 : 모델의 입력 셀에는 수치를, 결과 셀에는 방정식 및 관수를 기입하여 주십시오. 이것의 가정과 예측의 호출은, 이미 시뮬레이션 프로파일 존재하지 않으면 않습니다. 모델에 입력 가정의 기입을 행하는 것은 다음의 순서를 행하여 주십시오.

- 시뮬레이션 프로파일 있는 것을 확인하여 주십시오. 또는 신규 프로파일을 작성하거나 (리스크 시뮬레이터 (*Risk Simulator*) / 신규 시뮬레이션 프로파일 작성 (*New Simulation Profile*), 이미 보존 되어 있는 프로파일을 열어 주십시오.
- 가정을 기입하고자 하는 셀을 선택하여 주십시오 (예 : 기초 시뮬레이션 모델의 셀 G8 을 참조 하십시오).
- 리스크 시뮬레이터 (*Risk Simulator*) / 입력 가정의 설정 (*Set Input Assumption*) 을 클릭하거나, 또는 리스크 시뮬레이터의 톨의 세계의 아이콘을 클릭하여 주십시오.
- 관련이 있는 분포법을 선택하고, 분포의 파라미터를 입력한 후에, 모델의 입력 가정을 지정하기 위하여 **OK**를 클릭하여 주십시오. (Figure 2.3 참조)

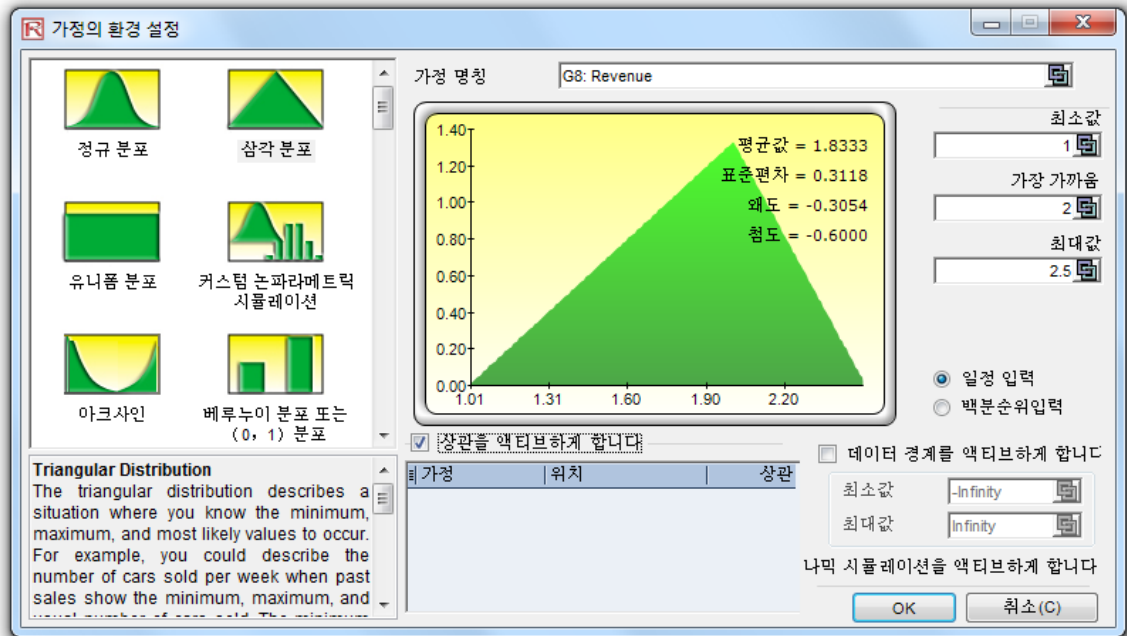


Figure 2.3 – 입력 가정의 설정

가정의 환경 설정에는 시뮬레이션에 몇 개의 중요한 항목이 있습니다. Figure 2.4 는 그러한 항목이 표시되어 있습니다.

- ② **가정 명칭** : 기입된 가정의 구별이 되도록, 각 가정의 명칭을 붙이고 있습니다. 짧고 알기 쉽게 명칭을 붙이는 것이 가장 적절하겠습니다.
- ② **분포 갤러리** : 좌측에는 소프트웨어에 내장되어 있는 여러 가지 분포법이 표시되어 있습니다. 분포의 화면 표시를 변경하는 것은 분포 갤러리를 열거 커다란 아이콘 및 작은 아이콘을 우클릭하여 보시기 바랍니다. 42 개의 분포법이 표시됩니다.
- ② **입력 파라미터** : 분포법의 선택에 의하여, 관련되는 파라미터가 표시됩니다. 덧붙여 파라미터는 직접 입력이 됩니다, 워크 시트의 특유의 셀에 파라미터를 링크시키는 것도 가능합니다. 가정 파라미터가 변동 하지 않고 알수 있을 때 이 툴을 사용하면 편리합니다. 파라미터를 보이게 표시하고자 할 때에, 또는 실행중에 파라미터를 변동하고자 할 때에 워크 시트의 셀에 링크 시키면 사용이 보다 충실합니다. (☐의 아이콘을 링크시켜 입력 파라미터를 워크 시트 셀에 링크 시키시기 바랍니다.)
- ② **가능한 데이터 환경** : 이 툴은 대부분 평균적인 분석자가 사용하지는 않습니다만, 분포 가정을 구분하는 것에 사용합니다. 예를 들어, 정규 분포를 선택했다고 하면, 이론상, 경계는 마이너스 무한과 플러스 무한의 사이에 있다고 합니다. 그러나, 실천에서는 시뮬레이션 된 변수는 좁은 범위내에서 존재하고, 이 범위는 적절한 분포를 구분하기 위하여 기입됩니다.
- ② **상관** : Pair 의 관계에 있는 상관은 입력 가정으로 지정하는 것이 가능합니다. 가정이 필요한 경우, 상관의 체크 박스를 클릭하는 것을 잊지 마십시오. 이 옵션은 **리스크 시뮬레이터 (Risk Simulator) | 시뮬레이션**

프로필의 편집 (Edit Simulation Profile) 을 클릭하면 표시됩니다. 상관의 지정과 상관이 불러오는 모델에의 효과 상세는、이 장의 최후에 있는 상관에 대하여 이론을 참조 하시기 바랍니다. 또、분포의 구분 및、또는 다른 가정에 분포를 관련 시키는 것이 가능합니다만、양방이 동시에 실행되지 않는 것에 주의하시기 바랍니다.

- ㉔ **간단한 기술** : 갤러리의 각 분포의 기술이 표시되어 있습니다. 필요한 입력 파라미터와 선택된 분포의 최적의 사용법 및 케이스가 기술됩니다. 소프트웨어에 내장되어 있는 각 분포의 상세는 몬테카를로 · 시뮬레이션을 위한 확률 분포의 장을 참조 하십시오.

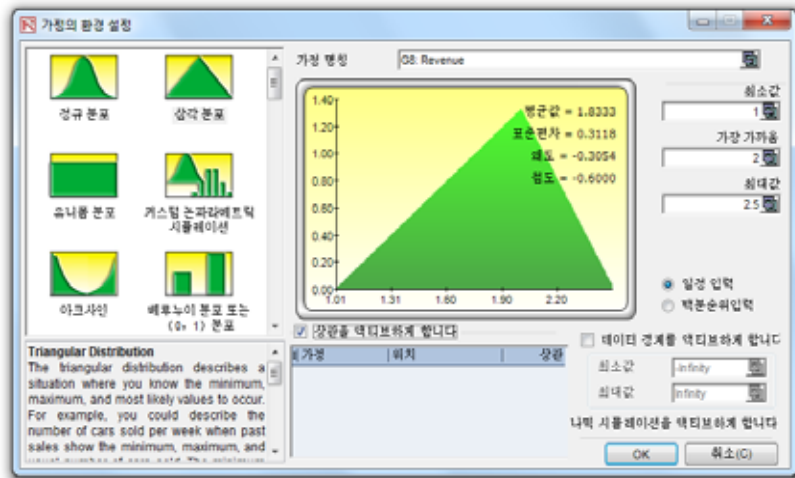


Figure 2.4 – 가정의 환경 설정

메모 : 예증을 시행할 경우에 다른 가정의 설정을 G9 에서 행하여 주십시오. 이 때에 정규 분포를 선택하고、최소값은 0.9, 최고값은 1.1 을 기입하여 주십시오. 그 후에、예측 결과의 설정을 다음 스텝과 같이 행하여 주십시오.

3.

다음 단계로서 모델의 예측 결과를 설정하십시오. 출력 셀은 방정식에 관수를 입력하십시오. 예측 결과의 설정 프로세스는 다음과 같습니다.

- 가정을 설정한 셀을 선택하여 주십시오 (예를 들어、기초 시뮬레이션 모델 에에서는、셀 G10) .
- 리스크 · 시뮬레이터 (Risk Simulator)** 를 선택하고、예측 결과의 설정 (**Set Output Forecast**) 를 클릭하여 주십시오. 또는 리스크 시뮬레이터의 툴바의 네 번째의 아이콘을 클릭하여 주십시오 (Figure 1.3 참조) .
- 필요한 정보를 기입하고、**OK**를 클릭하여 주십시오.

Figure 2.5 는 예측 결과의 환경 설정을 표시하고 있습니다.

- ㉔ **예측 결과의 명칭** : 예측 결과의 셀 명칭. 커다란 규모의 모델에서 복수의 예측 결과 셀을 설정할 때에 명칭을 구별하는 것으로、보다 보기 쉽게、보다 빠르게 결과를 구하는 것이 가능합니다. 이 단순한 스텝의 중요성을 잊지 마십시오. 짧고 알기 쉽게 명칭을 붙이는 것이 가장 적절합니다.

- ② **예측의 정도** : 시뮬레이션으로 어느 정도의 시행 수가 필요한가 감으로 맞추는 대신에, 정도와 에러 컨트롤의 설정을 사용하여 나누는 것이 가능합니다. 시뮬레이션 시행 회수를 자동 과정으로 하는 것으로, 에러와 정도의 조합의 시뮬레이션이 계산적으로 필요 시행 회수를 달성하면, 시뮬레이션은 정지하고, 필요한 정도의 달성을 연결하는 것으로, 필요한 시뮬레이션 회수를 사정에 추측할 필요가 없습니다. 상세는 에러 컨트롤과 정도의 장을 참조하십시오.
- ② **예측 결과의 화면의 표시** : 사용자의 희망으로 특정의 예측 화면의 표시, 비표시가 설정됩니다. 디폴트로, 항상 예측 차트가 표시되고 있습니다.

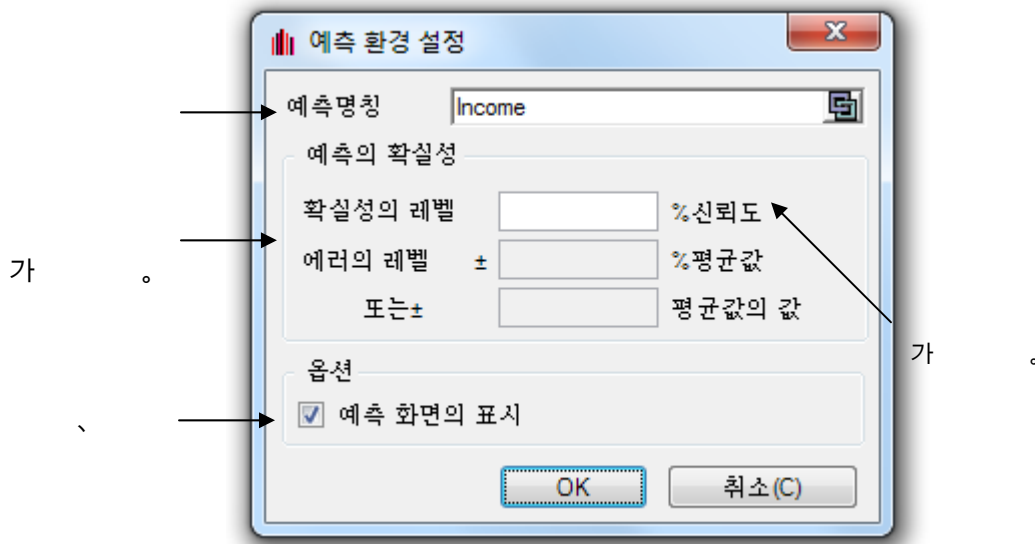


Figure 2.5 – 예측 결과의 설정

4. 시뮬레이션 실행

리스크 시뮬레이터 (Risk Simulator) / 시뮬레이션의 실행 (Run Simulation) 을 클릭하거나 **실행 (Run)** 아이콘을 클릭하면 (리스크 시뮬레이터 툴의 여덟개째의 아이콘) 시뮬레이션을 실행합니다. 재실행하는 경우, 또는 실행을 중단하는 경우는, 시뮬레이션을 초기화하는 것을 잊지 마십시오. 이를 위하여 **리스크 시뮬레이터 (Risk Simulator) / 시뮬레이터의 초기화 (Reset Simulation)** 를 클릭하고 시뮬레이션 툴 바의 10 번째의 아이콘을 클릭하여 주십시오. 즉 스텝 기능 (**리스크 시뮬레이터 (Risk Simulator) / 스텝 시뮬레이션 (Step Simulation)** 또는 툴 바의 아홉번째의 아이콘) 은 한번에 한번 시행하는 것이 가능합니다. 시뮬레이션의 교재로서 사용하면 적합할 것입니다. (각 시행에서 셀에 표시되 가정값이 매회 변화해 가면, 모델의 계산이 매회 행하여 지는 것이 증명됩니다)。

5. 예측 결과의 해석

몬테카를로·시뮬레이션의 최후 단계로서 예측 결과의 차트의 해석이 있습니다. Figures 2.6 에서 13 까지 예측 결과 차트와 시뮬레이션 실행 후에 생성된 통계가 표시

됩니다. 통상, 다음의 제 항목은 시뮬레이션의 예측 결과의 해석에 필요하다고 합니다.

- ㉠ **예측 차트** : Figure 2.6 에서 참조 된 예측 차트는 시뮬레이션의 총 시행수에서 일어나 값의 도수의 확률 히스토그램입니다. 세로축은 총 시행수에서 벗어나 일어난 특정의 X 값의 도수를 표시하고, 한편, 누적 도수(Smooth 한 라인)은 예측의 실행에서 일어난 총확률의 모든 값, 그리고 X 값보다 밑의 값도 표시됩니다.
- ㉡ **예측 통계** : Figure 2.7 의 참조는 예측값의 분포를 4 개의 분포 단계로 정리하고 있습니다. 이것의 통계 평균값의 상세는 예측 통계 이해의 장을 참조하시기 바랍니다. 스페이스 바를 사용하면 통계도와 히스토그램의 표가 교대로 표시됩니다.

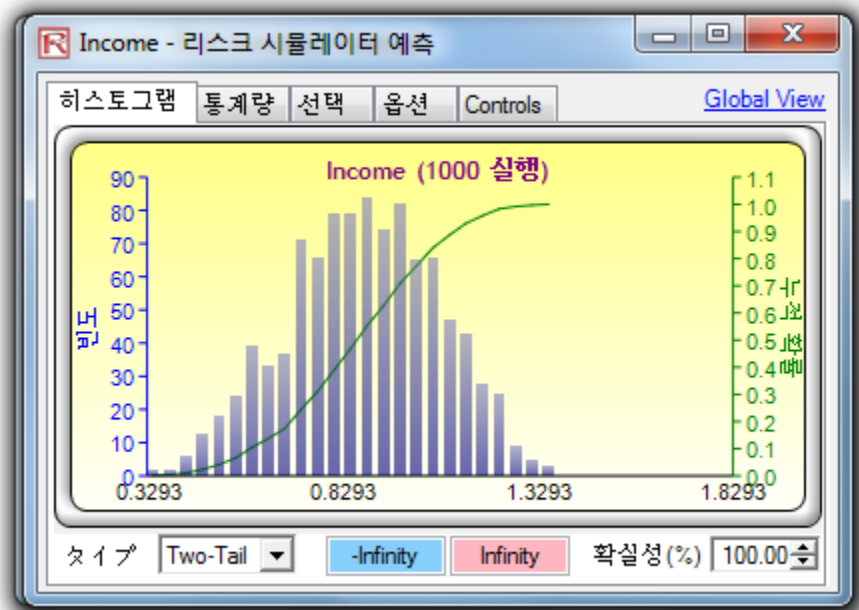


Figure 2.6 – 예측 차트

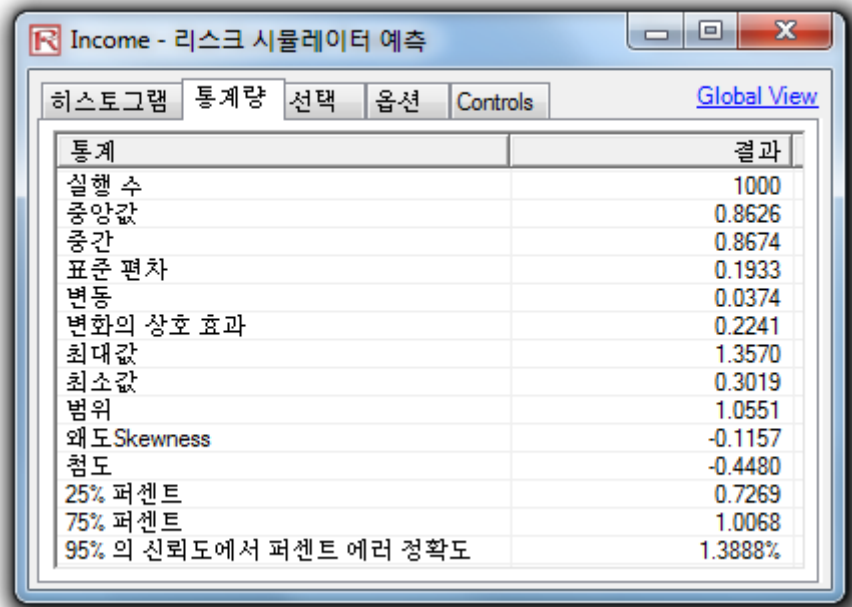


Figure 2.7 - 예측 통계

- ① 환경 설정 : 예측 차트의 환경 설정은 차트의 표시를 바꾸어 줍니다. 예를 들어, 항상 앞에 두는 *(Always On Top)* 선택
- ② 을 선택하면 예측 차트와 다른 프로그램을 사용하고 있어도 관계없이 항상 표시됩니다. 히스토그램 해결은 히스토그램의 bins 값을 5 bins 에서 100 bins 까지 변경하는 것이 가능합니다. 즉 *데이터 업데이트 옵션 (Data Update)* 은 시뮬레이션의 실행의 속도 vs 어느 정도의 기간에서 예측 차트가 업데이트하는 가를 설정할 수 있습니다. 즉, 각 시행에서의 예측 차트의 업데이트를 희망하면, 시뮬레이션의 실행 속도가 떨어집니다. 단, 이것은 사용자의 희망으로 설정하는 옵션으로, 시뮬레이션의 결과 (실행의 속도 이외) 에는 영향이 없습니다. 시뮬레이션의 실행 속도의 효율을 높이기 위하여, 시뮬레이션 실행중, Excel 의 화면을 축소하는 것이 가능합니다. 화면을 표시하기 위하여 메모리가 사용되지 않게 하기 위하여, 시뮬레이션의 실행 속도가 올라간다는 것입니다. *모든 것을 삭제 (Clear All)* 와 *모든 축소 (Minimize All)* 는 열려 있는 모든 예측 차트를 컨트롤합니다.

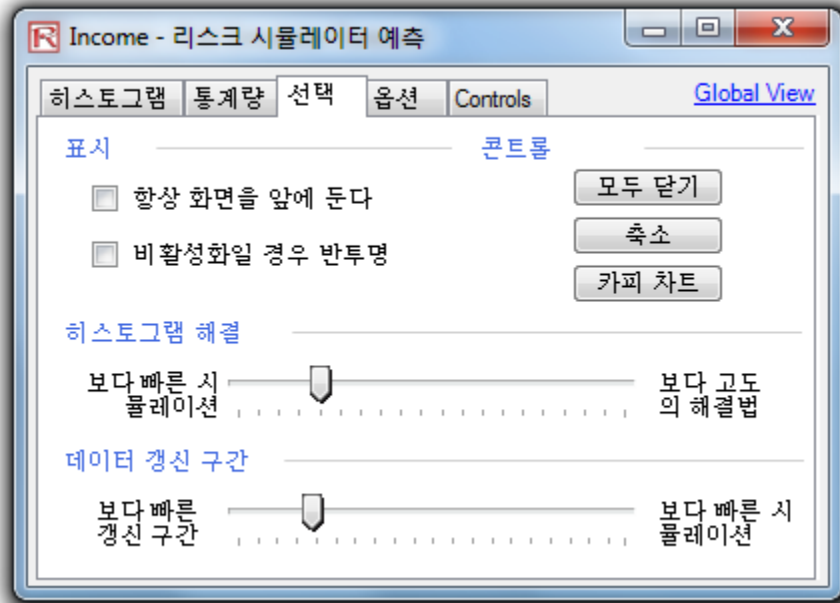
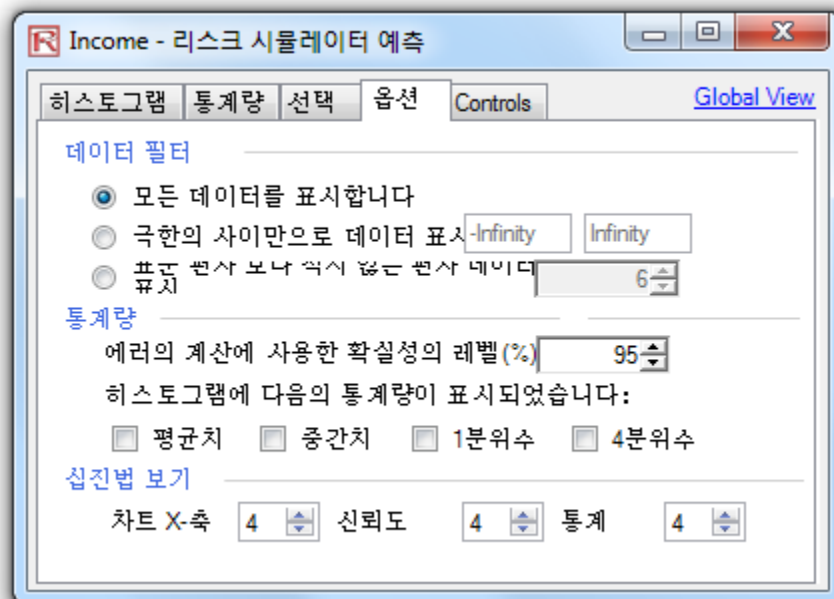


Figure 2.8 – 예측 차트의 환경 설정

- ② **옵션** : 이 예측 차트 옵션은 모든 예측 데이터를 표시하기 위하여, 또는 지정된 구간, 표준 편차에 해당 하는 입력·출력값을 필터링하기 위함입니다. 즉 정도 레벨을 여기에 지정하여, 통계표에 여러 레벨을 산출하는 것도 가능합니다. 상세한 정도와 여러 컨트롤이 장을 보시기 바랍니다. 다음의 통계를 표시하는 (*Show the Following Statistics*) *사용자의 자유로운 환경설정*에 있어 평균값, 중간값, 최초의 4 분위수와 네번째 4 분위수계열(25 번째와 75 번째 percentiles(백분 순위))가 예측 차트에 표시되지 않으면 않습니다.



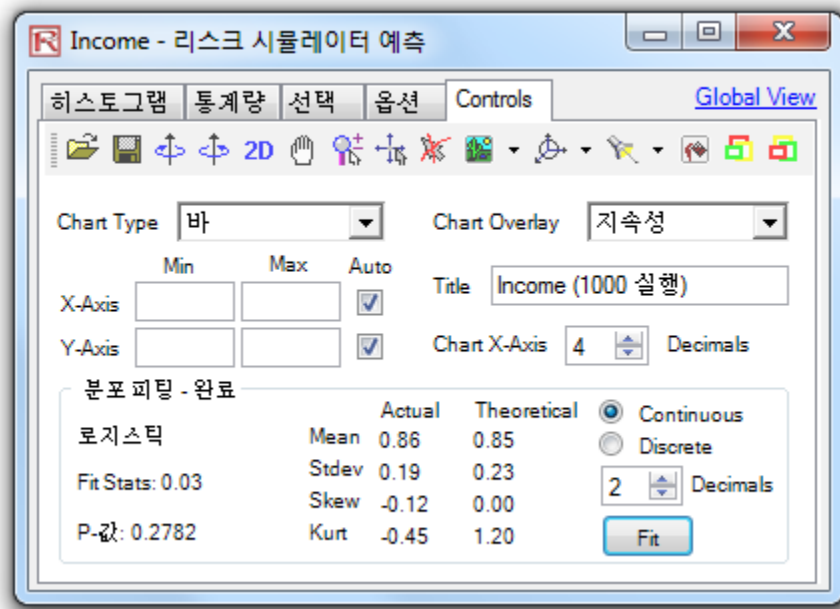


Figure 2.9 - 예측 차트 옵션

예측 차트와 신뢰 구간의 사용에 대하여 예측 차트에서는 신뢰 구간이라 불리우는 발생 확률을 지정할 수 있습니다. 즉, 두개의 값이 부여된 경우에, 어느 정도의 확률에서 결과값이 이 두 값의 사이에서 발생하는가 하는 것이 문제입니다. Figure 2.10 에서는 90%의 확률로 최종 결과 (이 케이스에서는 수입 레벨) 은\$0.5307 과 \$1.1739 의 구간에서 받아들여 지도록 표시하고 있습니다. 양측 신뢰 구간은 먼저, 양측 타입을 선택하고, 필요한 정도값을 기입하고 (예를 들어 : 90) TAB 키를 입력하여 주십시오. 정도값에 적합한 두 개의 값이 표시됩니다. 예증에서는 5%의 확률로 수입이\$0.5307 을 밑 돌고, 5%의 확률로 수입이\$1.1739 를 상회하도록 표시되고 있습니다. 이것은 양측 신뢰 구간이 대칭적으로 중간 구간인가, 50 번째의 percentile (백분순위)값입니다. 어느 쪽이든, 양방은 같은 확률을 가지고 있습니다.

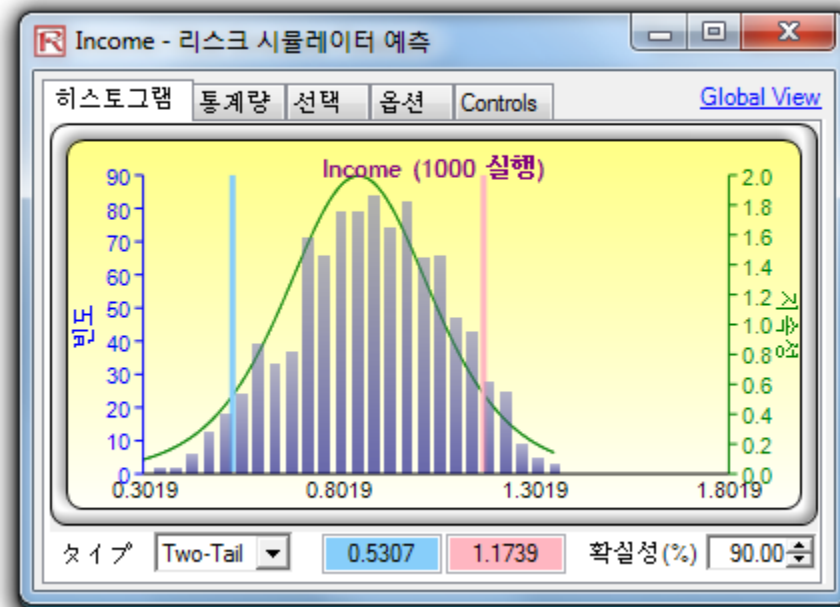


Figure 2.10 – 양측 신뢰 구간의 예측 차트

대신에 평균 확률의 계산이 가능합니다. Figure 2.11 에서는 좌측 선택을 95%의 신뢰도 (예 : 좌측을 선택하고, 정도의 레벨에 95 로 기입하고, TAB 키를 입력하도록 주시기 바랍니다.) 합니다. 이것은 95%의 확률로 수입이\$1.1739 를 밀돌거나, 5%의 확률로 수입이\$1.1739 를 상회한다는 것을 표시합니다. 이것은 Figure 2.10.로 표시되고 있는 결과에 정확히 기초한 수치입니다.

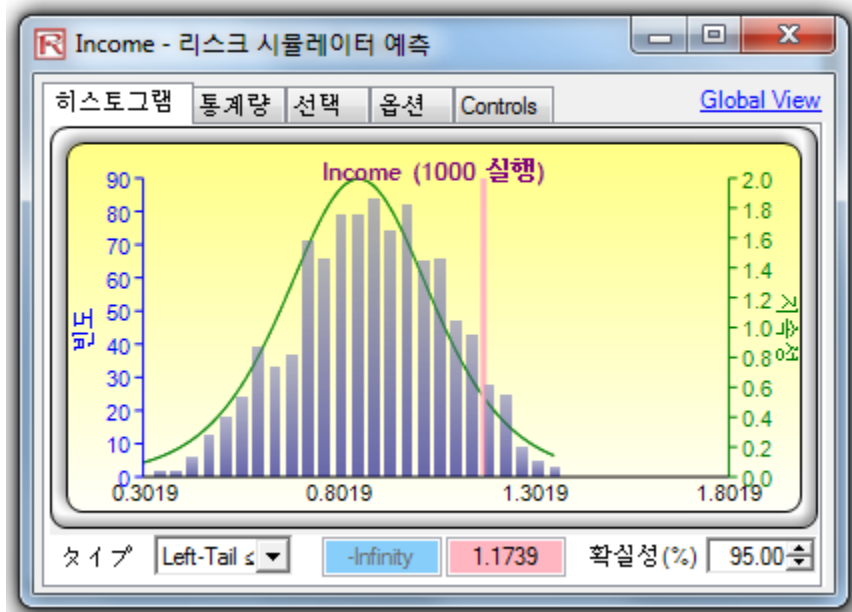


Figure 2.11 – 편측 신뢰 구간의 예측 차트

신뢰 구간에 대하여도 조금 덧붙이면 (정도 레벨을 부여하여 주고, 중요한 수입값을 보여 줍니다) 부여된 수입값의 정도도 계산할 수 있습니다. 예를 들어, 어느 정도의

확률로, 수입이\$1 를 밑도는 것일가가 문제입니다. 이것을 해결하는 것은 좌측의 확률 타입을 선택하고, 1의 수치를 기입하고, TAB 키를 클릭하면, 적합한 확률값이 표시됩니다. (이 경우에는 74.30%의 확률로 수입이 \$1 를 밑도는 결과가 나옵니다)。

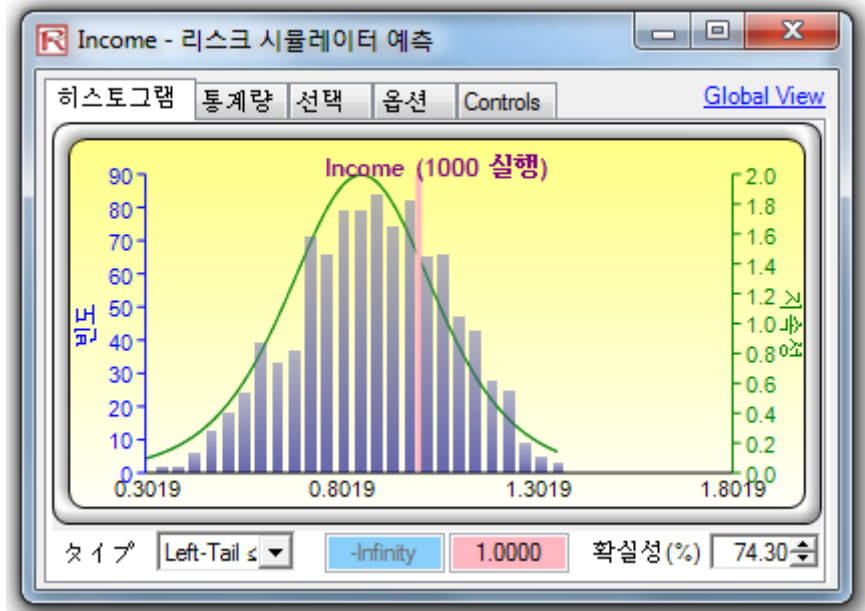


Figure 2.12 – 예측 차트 확률 평가

즉, 좌측의 확률 타입을 선택하고, 1의 값을 입력 박스에 기입하고, TAB 키를 클릭하여 주십시오. 좌측의 확률은 값 1을 넘는 것을 표시합니다. 즉, 수입이\$1 를 넘을 확률입니다 (이 케이스에서는 25.70%의 확률로 수입이\$1 를 넘을 것을 알 수 있습니다)。

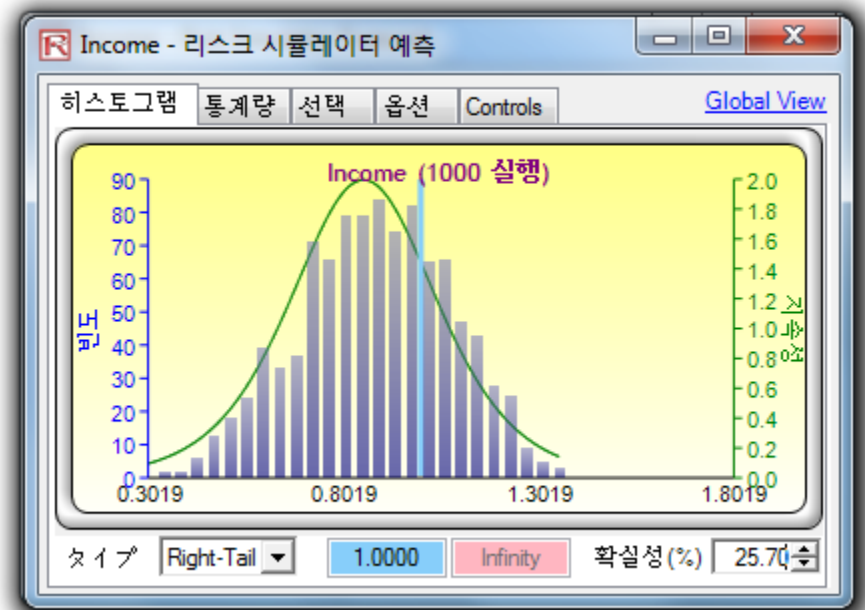


Figure 2.13 – 예측 차트의 확률 평가

어드바이스 : 예측 화면의 우측 끝을 클릭하여, 드랙하는 것으로 사이즈 변환이 가능합니다. 최후에 시뮬레이션을 재실행하기 전에 매회 Reset 하는 것을 추천합니다.

Reset 에서는 다음의 메뉴를 클릭하여 주십시오 (**리스크 시뮬레이터 (Risk Simulator) / 시뮬레이션의 Reset (Reset Simulation)**) 。 확실한 값 및 좌우에 값을 기입할 때에는, 차트와 결과를 업데이트 하는 것에 TAB 키를 누르는 것을 잊지 마십시오.

상관의 기초

두개의 변수 사이의 관계의 강함과 방향은 상관 계수의 측정으로, -1.0 과 +1.0 의 사이의 값에 해당됩니다. 즉, 상관 관계 부호 (두 개의 변수 사이의 관계가 플러스 인가, 마이너스 인가) 와, 관계의 크기와 강도 (높이는 상관 계수의 절대값, 강도와 관계) 에 의하여 분해됩니다. 상관 계수는 여러 가지 방법으로 계산 할 수 있습니다. 처음에, 두 개의 변수 x 와 y 를 사용한 상관 r 을 수동으로 계산합니다.

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

2 번째의 방법으로, Excel 의 CORREL 기능을 사용하는 것입니다. 예를 들어, x 와 y 의데이터 포인트가 A1:B10 의 셀에 기입되는 경우에, CORREL (A1:A10, B1:B10)을 표시합니다.

3 번째의 방법으로서, 리스크 · 시뮬레이터의 멀티 · 핏 툴 (Multi-Fit Tool) 을 실행하면, 결과로서 일어난 상관 매트릭스가 계산되고 표시됩니다.

상관관계는 인과 관계를 의미하지 않는 것에 주목하여 주십시오. 두 개의 무관계의 확률 변수는, 어떤 상관 관계를 지정하지 않으면 않습니다, 원인은 합의되지 않았습니다 (예를 들어, 태양의 흑점의 활동 주식 시장에 있어서 이벤트는 상관되어 있습니다만, 2 개의 사이에, 인과 관계가 전혀 없습니다) 。

2 개의 일반적 타입의 상관이 있습니다 : 파라메트릭과 논 파라메트릭 상관입니다. 피아슨 상관관계법은, 무엇보다도 일반적인 상관 측정법이 있어, 통상, 간단히 상관관계로서 표시하고 있습니다. 그러나, 피아슨 상관법은 파라메트릭법이 있는 것에서, 상관관계에 있는 서로의 변수는 기본적으로 정규 분포에서, 변수 사이의 관계는 선형이 되지 않으면 안되는 것을 표시하고 있습니다. 몬테카를로 시뮬레이션에서 잘 볼 수 있는, 이것의 조건이 만족되지 않는 케이스에서는, 논 파라메트릭의 대응 측이 무엇보다도 중요합니다. 스피아만의 순위 상관과 겐틀의 다우의 2 개가 대책안입니다.

스피아만 상관은 무엇보다도 일반적으로 사용되고, 몬테카를로 · 시뮬레이션의 환경에 적용하는 것이 무엇보다도 적절합니다. 선형성 혹은 정규 분포상에 의존하지 않기 위하여, 다른 분포를 가지고 있는 변수간의 상관이 적용되는 것을

표시합니다. 스피아맨의 상관을 계산하기 위하여는, 먼저, 모든 x 와 y 의 변수를 순위별로 랭크시키고 피아슨 상관의 계산을 행합니다.

리스크 시뮬레이터의 경우, 사용된 상관은, 보다 완강한 논파라메트릭스 스피아맨의 랭크 상관입니다. 단, 시뮬레이션의 과정을 간소화하여, Excel 의 상관 기능과 일성관을 가지기 위하여는, 필요되는 상관의 입력은, 피아슨의 상관계수입니다. 그 후에, 리스크 시뮬레이터는, 독특한 알고리즘을 적용하고, 과정의 간소화에 의하여, 이것을 스피아맨의 랭크를 교환합니다. 단, 사용자 인터페이스를 간소화하는 것은, 보다 일반적인 피아슨의 적용 상관의 추가 입력을 가능하게 합니다 (예, EXCEL 의 CORREL 기능을 사용한 계산) 。 한편, 수학적인 코드에서는, 이것의 심플한 상관을 분포 시뮬레이션을 위하여 기초인 스피아맨의 순위 상관으로 교환합니다.

리스크 시뮬레이션에서 상관을 적용하는 것에 대하여

상관은 리스크 시뮬레이터에 여러 가지 방법으로 적용할 수 있습니다:

- ① 가정을 정의 할 때 (*Risk Simulator*) | 입력 가정의 설정 (*Set Input Assumption*)), 분포 갤러리에서 상관 매트릭스 격자내에 상관을 기입하여 주십시오.
- ② 데이터의 존재가 있을 경우, 멀티 · 적합 툴에서(리스크 시뮬레이터(*Risk Simulator*) | 툴(*Tools*) | 분포 적합(*Distributional Fitting*) | 복수의 변수(*Multiple Variables*)), 분포 적합을 표시하고, 페어 와이즈 변수 사이의 상관 매트릭스를 얻기 위하여 실행합니다. 만약, 시뮬레이션 프로필이 존재한다면, 적합한 가정은 자동적으로 관련이 있는 유사관 값을 포함하고 있는 것으로 생각됩니다.
- ③ 데이터의 존재가 있을 때, 리스크 시뮬레이터(*Risk Simulator*) | 상관의 편집(*Edit Correlations*)을 클릭하고, 한 사람의 사용자 인터페이스내의 모든 사정의 페어 와이즈 상관의 기입을 행합니다.

상관 매트릭스가 플러스 직행렬이 아니면 않되는 것에 주목하시기 바랍니다. 이것은, 상관은 수학적으로 유효하지 않으면 않되는 것을 표시하고 있습니다. 예를 들어, 세 개 변수의 상관을 계산하고자 하고 있습니다. 어떤 특정의 년도의 대학원생 학년, 그들이 주에 소비하는 맥주의 수와, 그들이 주에 공부하는 시간 수. 다음과 같은 상관 관계의 존재가 가정됩니다.

학생과 맥주: - 맥주의 소비가 높을수록, 학년이 낮고 (시험에 결석)

학생과 공부 시간: + 공부 시간이 많을수록, 고학년이다.

맥주와 공부시간: - 맥주의 소비량이 높을수록, 공부 시간이 적고 (항상 놓고, 마시기)

단, 학년과 공부 사이의 상관 관계를 입력하고, 상관 계수가 높은 경우, 상관 매트릭스는 마이너스 직행렬이 되게 됩니다. 이것은, 이론에, 상관의 필요조건에, 그리고 수학적인 행렬에 대립하는 것이 됩니다. 단, 작은 상관은 때에 따라, 이론이 불충분한 경우에도 의연하게 기능합니다. 마이너스 직행렬의, 또는 틀린

상관 행렬을 기입하였을 경우에는, 리스크 시뮬레이터는 자동적으로 표시되고, 상관 관계 (동일의 상대적 강함과 동일의 부호) 의 전체적인 구성을 유지하면서도, 이것의 상관을 마이너스 직행력 근접하게 하기 위하여 조절합니다.

몬테카를로·시뮬레이션에서의 상관의 영향

시뮬레이션에서의 변수의 상관에 필요하다고 생각되는 계산은 복잡해도, 결과로서 표시되는 영향이 꽤 명확합니다. Figure 2.14 는, 단순 상관 모델을 예증 폴더 내의 상관 영향 모델)을 표시합니다. 수입을 위한 계산은, 단순히 가격을 양에 곱한 것입니다. 가격과 양의 상이의 제로 상관, 플러스 상관(+0.8)과 마이너스 상관(-0.8)에 같은 모델이 복사되어 있습니다.

	B	C	D	E	F	G	H
1							
2							
3			상관 모델				
4			Without Correlation	Positive Correlation	Negative Correlation		
5		Price	\$2.00	\$2.00	\$2.00		
6		Quantity	1.00	1.00	1.00		
7		Revenue	\$2.00	\$2.00	\$2.00		
8							
9		이 모델을 복사하기 위하여는 다음의 가정을 사용하십시오:					
10		Price는 Triangular 분포로 설정하고 (1.8, 2.0, 2.2) 종자값 123456으로					
11		상관관계가 0.0, +0.8, -0.8 일때 Quantity를 정규 분포 (0.9, 1.1)로 세팅하고 1,000번을 실행하십시오.					

Figure 2.14 – 단일 상관 모델

결과로서의 통계는, Figure 2.15 에 표시되어 있습니다. 상관이 포함되지 않은 모델의 표준 편차는, 0.1876 의 플러스의 상관과 0.0706 의 마이너스의 상관에 비교하면 0.1450 에 있는 것을 주목하여 주십시오. 이것은, 심플한 모델에서는, 마이너스의 상관은 분포의 평균적 분포를 감소하는 경향을 보여, 보다 크게 평균적 분포를 가진 플러스의 상관에 비교하여, 밀접하게 집중된 예측 분포를 작성하는 경향이 있습니다. 단, 평균값은 비교적 안정되게 남습니다. 이것은, 상관은 프로젝트의 기대치를 거의 변화시키지 않으나, 프로젝트의 리스크를 감소시키거나, 증가시키거나 하는 것을 의미합니다.



Figure 2.15 -상관의 결과

Figure 2.16 은, 시뮬레이션의 실행 후에 표시되는 결과로서, 가정의 미가공 데이터를 추출하고, 변수 사이의 상관을 계산합니다. 표는, 입력 가정이 시뮬레이션으로 재현되는 것을 표시합니다. 즉, 자신이 +0.9 과 -0.9 의 상관을 기입하면, 결과로서 표시된 시뮬레이션의 값은 상관과 같게 됩니다.

이하는 시뮬레이션에서 추출된 Raw values입니다. 이것들은 실제로 가정에 입력되는 상관을 분석하기 위하여 상관되고, 실제로 모델화 되는 상관이 됩니다. 피어슨 상관은 실행 파라메트 상관관계의 상호 요퍼이며, 결과는 입력된 상관(+0.80 과 -0.80)을 지시하는 것이 실제로 변수값사이의 상관입니다.
자세한 사항은 Dr. Johnathan Mun 의 "Modeling Risk, Second Edition" (Wiley 2010) 참조.

Price Positive Correlation	Quantity Positive Correlation		Price Negative Correlation	Quantity Negative Correlation	
1.95	0.91		1.89	1.06	
1.92	0.95		1.98	1.05	
2.02	1.04	Pearson's Correlation:	1.89	1.09	Pearson's Correlation:
2.04	1.03		1.88	1.04	
1.89	0.91	0.80	1.96	0.93	-0.80
1.98	1.05		2.02	0.93	
2.05	1.03		2.00	1.02	

Figure 2.16 -상관 재현

정밀도와 에러 컨트롤

몬테카를로 시뮬레이션의 기능 중 강력한 기능은 정밀도 컨트롤입니다. 예를 들어, 복잡한 모델에서 실행하는 것으로 어느 정도의 시행수가 충분합니까. 정밀도 컨트롤은 미리 설정한 정밀도 수준에 도달하는 경우 시뮬레이션을 정지시키는 것이 가능한 것으로 적절한 시행 수의 추정이라는 것으로, 적절한 시행수의 추정이라는 작업을 제거하고 있습니다.

정밀도의 컨트롤의 기능은, 예측을 어느 정도 확실히 하고싶은가에 대하여 스스로의 설정을 가능하게 해 줍니다. 일반적으로, 시행수가 많을수록, 신뢰 구간이 좁게 되어, 통계는 보다 확실하게 됩니다. 리스크 시뮬레이터에서 적용되고 있는 정밀도 컨트롤은, 신뢰 구간의 특성을 사용하고, 언제 통계가 있는 특성의 정밀도가 달성되는가를 보여 주고 있습니다. 각 예측의 확률 정밀도의 레벨에 특성의 신뢰 구간을 지정할 수 있습니다.

세계의 다른 언어를 혼동하지 않도록 주의하십시오: 에러, 정밀도와 신뢰입니다. 이것들이 비슷한 것 같지만, 각각 가지고 있는 개념은 확실히 다릅니다. 이것의 샘플은 다음에 표시되어 있습니다. 예를 들어, 타코스 (타코스의 껍질) 제조업자라면, 100 의 조가비가 들어 있는 상자 안에서 평균 어느 정도의 쉘이 깨져 있는가를 알고 싶어합니다. 하나의 방법으로는, 100 의 쉘이 들어가는 Prepackage 된 상자의 샘플을 주문하여, 열고, 어느 정도의 쉘이 불량인가를 세어봅니다.

하루에 1,000,000 상자 (이것이 모집단이 됩니다) 를 제조하고, 랜덤으로 그 중에서 10 상자를 열어 봅니다 (이것이 샘플 사이즈가 되어, 시뮬레이션의 시행 수로도 생각되고 있습니다) 。 열어 본 상자의 안에서 불량 쉘의 수는 다음과 같습니다 : 24, 22, 4, 15, 33, 32, 4, 1, 45, 와 2. 계산된 불량품의 쉘의 평균수는 18.2 입니다. 이것의 샘플, 또는 10 의 시행 회수에 기초한 평균값은 18.2Unit 인 것에 대하여, 샘플에 기초한 80%의 신뢰 구간은 2 에서 33 Unit 이 됩니다 (이것은, 실행되는 시행 회수, 또는 샘플 사이즈에 기초한 80%에 대하여, 결과로서 불량품의 수는 2 에서 33Unit 의 중에서 해당되는 것입니다) 。 단, 어떤식으로 하여 18.2 이 바른 평균값인가를 정할 수 있겠습니까? 10 의 시행 회수는, 결과의 정도를 정하는데 충분합니까? 2 에서 33 의 사이의 신뢰 구간에서, 지나치게 넓은, 변동적이라고 할 수 있습니다. 예를 들어 90%의 빈도 데이터 쉘의 에러의 도수가 ± 2 라고 하는 한층 정확한 평균값이 필요한 경우를 가정하면, 이것은, 하루에 제조된 모든 1,000,000

상자를 열면, 이것중 900,000 상자에는, 특성의 평균값에서 평균 ± 2 의 개수의 범위의 불량인 쉘이 들어 있다는 것을 의미합니다. 이 정도의 정밀도를 가진 결과를 얻기 위하여는 어느 정도의 시행 회수, 또는 샘플 사이즈가 필요하겠습니까? 여기에서는, 두 개의 타코스가 에러 레벨을 표시하고, 90%의 정밀도를 표시하고 있습니다. 충분한 시행회수가 실행되는 경우, 90%의 신뢰 구간은, 90%의 정밀도 레벨과 같게 되어, 90%의 빈도, 에러 (오차) , 또는 신뢰도가 ± 2 와 같이 한층 정확한 평균의 척도를 얻을 수 있습니다. 예증과 같이, 평균값이 20Unit 라면, 90%의 신뢰 구간은 18 에서 22Unit 사이에 있다고 합시다. 그리고 그 신뢰도 90%의 정밀도를 가지고, 1,000,000 상자를 열은 경우, 그 안의 900,000 상자에는 18 에서 22Unit 의 불량한 타코스 쉘이 들어 있는 것이 됩니다. 이 정도는 충분한 시행 회수를 얻는

것에는, 다음의 샘플링 에러의 방식 $\bar{x} \pm Z \frac{s}{\sqrt{n}}$ 에 기초하지 않으면 안됩니다. $Z \frac{s}{\sqrt{n}}$ 는 타코스 2 개의 에러의 도수를 표시하고, $\bar{x}$ 는, 샘플의 평균값, Z 는, 90%의 정도

레벨의 표준 정규 Z-스코어가 됩니다. 또한 s 는, 표준 편차를 표시하고, n 은, 지정된 정도에 대비한 에러의 레벨을 얻기 위하여 필요한 시행 회수가 됩니다. Figures 2.17 과 2.18 는, 리스크 시뮬레이터로, 정밀도 컨트롤이 어떤 식으로 하여, 복수의 예측에 적용할 수 있는가를 표시합니다. 이것은, 사용자가 시뮬레이션을 실행하는 것이 어느 정도의 시행 회수가 필요한 가를 정하는 과정을 심플하게 하고 있습니다.

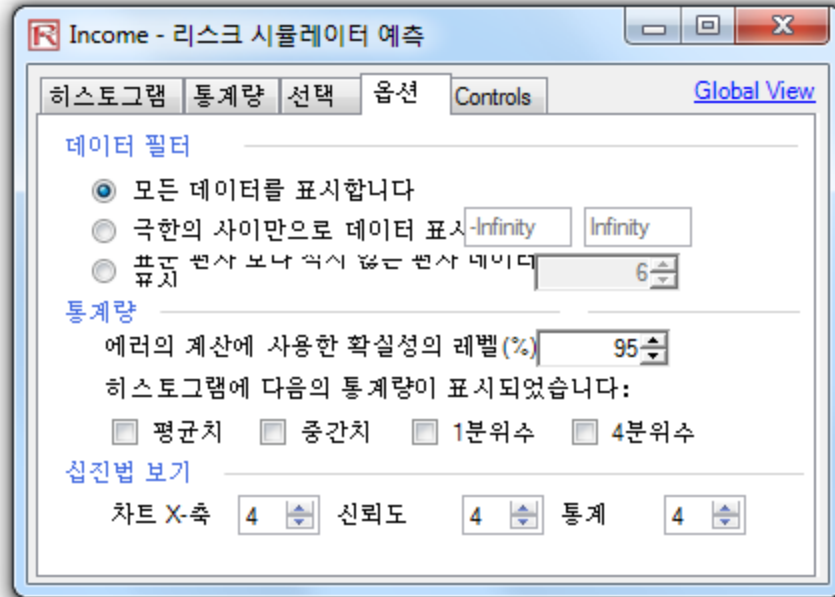


Figure 2.17 - 예측 정밀도 레벨의 설정

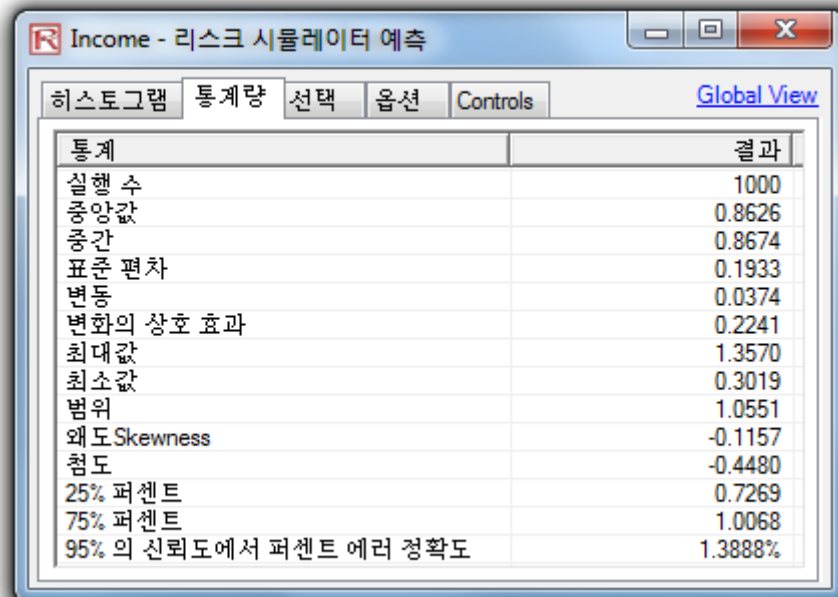


Figure 2.18 - 에러의 계산

예측 통계를 해독하는 것에 있어서 거의 모든 분포는 4 개의 모멘트까지 정의 가능합니다. 제 1 차 모멘트는, 분포의 위치, 또는 중심의 경향 (기대 리턴) 를 표시하고, 제 2 차 모멘트는, 이것의 폭, 또는 확장 (리스크) 를 표시하고, 제 3 차 모멘트는 비뚤어진 방향 (무엇보다도 확률이 높은 이벤트) 를 표시하고, 제 4 차 모멘트는, 우도, 또는 테이블의 두께 (파국적 손실, 또는 이익) 을 표시하고 있습니다. 이것의 모든 4 차 모멘트는, 분석 하에 프로젝트의 무엇보다도 알기 쉬운 시점을 얻기 위하여, 실행 시 계산 되어 해석되지 않으면 않습니다. 리스크 시뮬레이터는, 이것의 모든 4 차원 모멘트의 결과를 예측 차트의 통계의 표시에 부여합니다.

분포의 중심을 측정하는 것에 있어서 -제 1 차 모멘트

분포의 제 1 차 모멘트는, 어떤 특정의 프로젝트의 리턴의 예상 비율을 측정합니다. 즉, 프로젝트의 시나리오의 위치를 측정하고, 평균치로서의 확률로서 생각되는 결과를 측정합니다. 제 1 차 모멘트를 위하여 무엇보다도 일반적인 통계는, 평균 (평균값) , 메디안 (분포의 중심) 과 모드 (무엇보다도 일반적으로 발생하는 값) 이 포함되어 있습니다. Figure 2.19 에서는, 제 1 차 모멘트가 표시되어 있어, 이 케이스에서의 분포의 제 1 차 모멘트는 , 평균, 또는 평균치로 측정되고 있습니다.

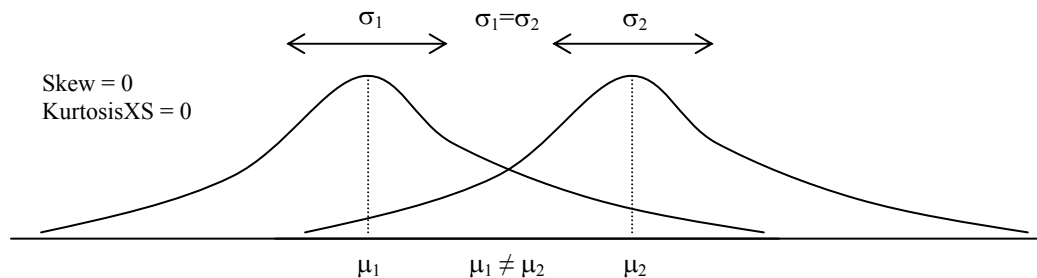


Figure 2.19 - 제 1 차 모멘트

분포의 파급의 측정 -제 2 차 모멘트

제 2 차 모멘트는, 분포의 파급, 즉 리스크를 측정합니다. 분포의 파급, 또는 분포의 폭은, 변수의 변동을 측정합니다. 이것은, 분포의 여러 가지 범위에 변수가 떨어지는 잠재성으로, 즉 결과의 시나리오의 잠재성입니다. Figure 2.20 는, 같은 제 1 차 모멘트 (같은 평균) 을 가진 2 개의 분포를 표시하고 있습니다만, 전혀 다른 제 2 차 모멘트, 또는 리스크를 표시하고 있습니다. 시각적으로는 Figure 2.21 에 무엇보다도 명확히 됩니다. 예증과 같이, 2 개의 주식이 있다고 하고, 작은 변동의 최초의 주식의 변동 (흑선에서 표시되고 있습니다) 에 대하여, 2 번째의 커다란 가격의 주식의 변동 (점선에서 표시되고 있습니다) 으로서 비교되고 있습니다. 한층 리스크가 높은 주식의 결과는 비교적 리스크가 낮은 주식에 비하여 알려져 있지 않음으로, 분명히, 투자자는 변동의 커다란 주식을 한층 높은 리스크로 보는 경향이 있습니다. Figure 2.21 의 종축은, 주식의 가격을 측정합니다만, 리스크가 높은

주식일수록, 잠재적인 결과의 폭넓은 범위를 표시하고 있습니다. 이 범위는, 분포의 폭 (횡축) 으로서 Figure 2.20 으로 해석되고 있어, 폭넓은 분포는 리스크의 높은 자산을 표시하고 있습니다. 따라서, 분포의 폭, 또는 범위는 변수의 리스크를 측정합니다.

Figure 2.20 의 양방의 분포는, 동일한 제 1 차 모멘트, 또는 중심의 경향을 가지고 있습니다만, 분명히 다른 분포에 있는 것에 주목하여 주십시오. 그 차이는, 분포의 폭에서 측정할 수 있습니다. 수학적, 그리고 통계적으로는, 변수의 폭, 또는 리스크는, 범위, 표준 편차(σ), 분산, 변동의 계수와 백분위 수를 포함한 여러 가지 통계를 통하여 측정할 수 있습니다.

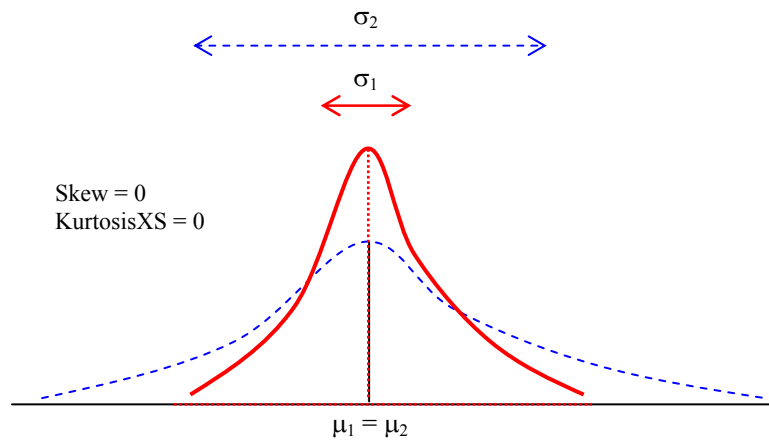


Figure 2.20 - 제 2 차 모멘트

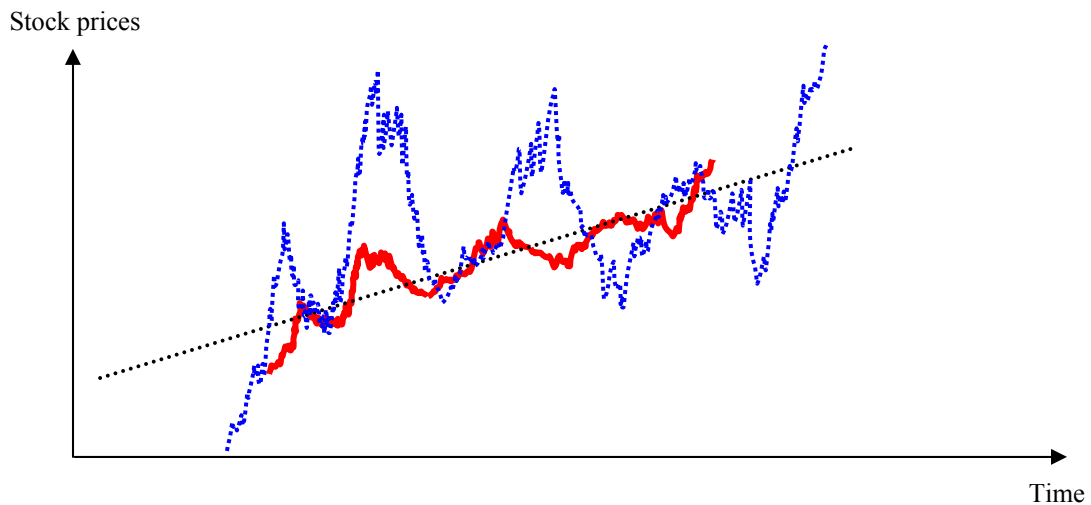


Figure 2.21 - 주가의 변동

분포의 왜도를 측정하는 것에 있어서—제 3 차 모멘트

제 3 차 모멘트는, 어떻게 하여 분포가 편측, 또는 역방면에서 밀리는가, 분포의 왜곡을 측정합니다. Figure 2.22 는, 마이너스, 및 좌측의 경사 (분포의 포인트의 꼬리를 좌측으로) 을 표시하고 있습니다. 또는, Figure 2.23 는, 플러스, 및 우측의 경사 (분포의 포인트의 꼬리를 우측으로) 을 표시하고 있습니다. 평균은, 항상 분포의 꼬리의 쪽이 기울어져 있으나, 중앙값은 항상 일정하게 유지 되고 있음으로. 이것을 다르게 보면, 평균은 움직이나, 표준 편차, 분산, 또는, 폭은 일정하게 머물러 있는 것이 됩니다. 만약 제 3 차 모멘트를 고려하지 않고, 기대 리턴 (예, 중간, 및 평균) 과 리스크 (표준 편차) 만을 관찰한 경우, 플러스의 경사가 있는 프로젝트가 잘못 선택됩니다. 예를 들어, 횡축이, 프로젝트의 순 수입을 표시하고 있는 경우, 분명히 플러스, 및 좌측 경사의 분포는 높은 확률에 높은 리턴(Figure 2.22)이 있어, 높은 확률로 낮은 리턴(Figure 2.23)에 비교하여 좋아할 것입니다.

따라서, 왜곡된 분포에서는, 중심값 (중위) 는 리턴의 보다 좋은 측정 척도로써, Figures 2.22 와 2.23 의 양방에서는 중심값이 동일하고, 리스크도 동일함으로, 순이익이 마이너스의 경사를 가진 분포의 프로젝트를 선택하는 것이 희망적이 됩니다. 프로젝트의 분포의 왜곡의 고려의 실패는, 바르지 않은 프로젝트를 선택한 것을 표시합니다 (예를 들어, 두 개의 프로젝트가 같은 제 1 차, 제 2 차 모멘트 등을 가지고 있는 경우, 이것은 양방과 같이 같은 리턴과 리스크의 프로필을 가지고 있는 것이 됨으로, 이것의 분포의 왜곡은 전혀 다를지도 모릅니다.

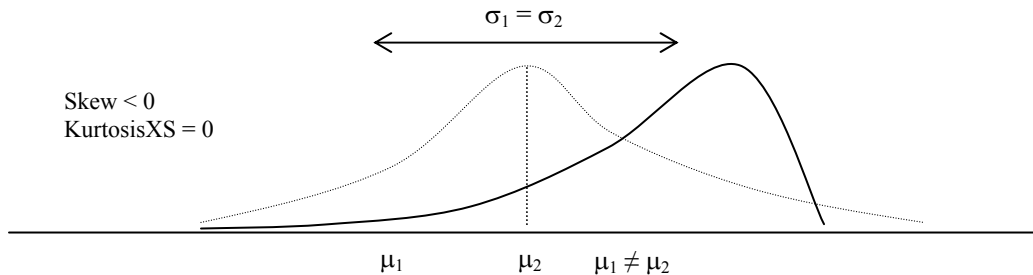


Figure 2.22 – 제 3 차 모멘트 (좌측 경사)

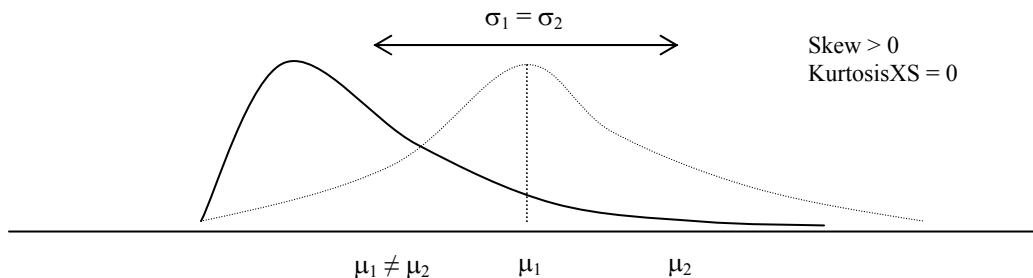


Figure 2.23 – 제 3 차 모멘트 (좌측 경사)

분포에서의 파국적인 Tail Event 를 측정하는 것에 있어서 —제 4 차 모멘트

제 4 차 모멘트, 또는 우도는, 분포의 피크를 측정합니다. Figure 2.24 는, 이 효과를 표시하고 있습니다. 백 그라운드(점선으로 표시되고 있습니다)는, 3.0 의 우도 혹은 0.0 의 초과분의 우도(KurtosisXS)의 정규 분포입니다. 리스크 시뮬레이터의 결과는, 우도의 평상 레벨로서 0 을 사용한 KurtosisXS 의 값을 표시합니다. 이것은, 플러스 값이 돌출한 꼬리 (학생의 T, 혹은 대수 정규 분포와 같은 급첨적 분포) 를 표시하고 있으나, 마이너스 KurtosisXS 는, 평평한 꼬리 (보통 분포와 같은 완첨적 분포) 를 표시하고 있는 것이됩니다. 굵은 선에서 그려진 분포는 높은 첨도를 가지고 있음으로, 곡서의 밑의 범위는 꼬리에 대하여 보다 두껍고, 중심부의 범위는 보다 얇게 됩니다. 이 조건은, Figure 2.24 에 있는 두 개의 분포와 같은 리스크 분석에 의하여 커다란 영향을 부여하고, 처음의 제 3 차 모멘트 (평균, 표준 편차와 왜도) 가 동일하다고 하면, 제 4 차 모멘트 (첨도) 는 다릅니다. 이 조건은, 리턴과 리스크가 같다고 하는 경우도, 극단, 및 파국적인 이벤트 (커다란 이익, 또는 커다란 손실) 의 발생 확률은, 높은 돌출의 분포 (예, 시장의 주식의 리턴은 급돌적, 또는 높은 첨도) 쪽이 커다란 것을 표시하고 있습니다. 프로젝트의 첨도를 무시하는 것은 유해하게 됩니다. 크게 낮거나, 초과한 첨도의 값은, 다운 사이즈의 리스크는 높은 것을 표시하고 있습니다(예, 프로젝트의 Value at Risk 의 값은 유의하지 않으면 않됩니다).

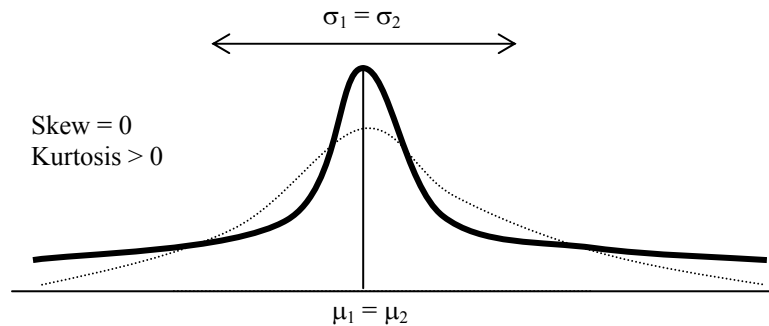


Figure 2.24 – 4

3.

예측이란, 미래의 예언을 하는 행동에 있어서, 이력적 데이터, 또는, 이력적 데이터가 존재하지 않는 경우는, 미래의 추측에 기초하여 행하여 집니다. 이력이 실재하는 경우, 정량적 혹은 통계적 어프로치가 최적입니다만, 이력적 데이터가 없는 경우, 가능성으로서는 정성적 혹은 판단의 적용이 적절하게 되어, 단 하나 방법이 됩니다. Figure 3.1 는, 무엇보다 일반적인 예측 방법을 표시하고 있습니다.

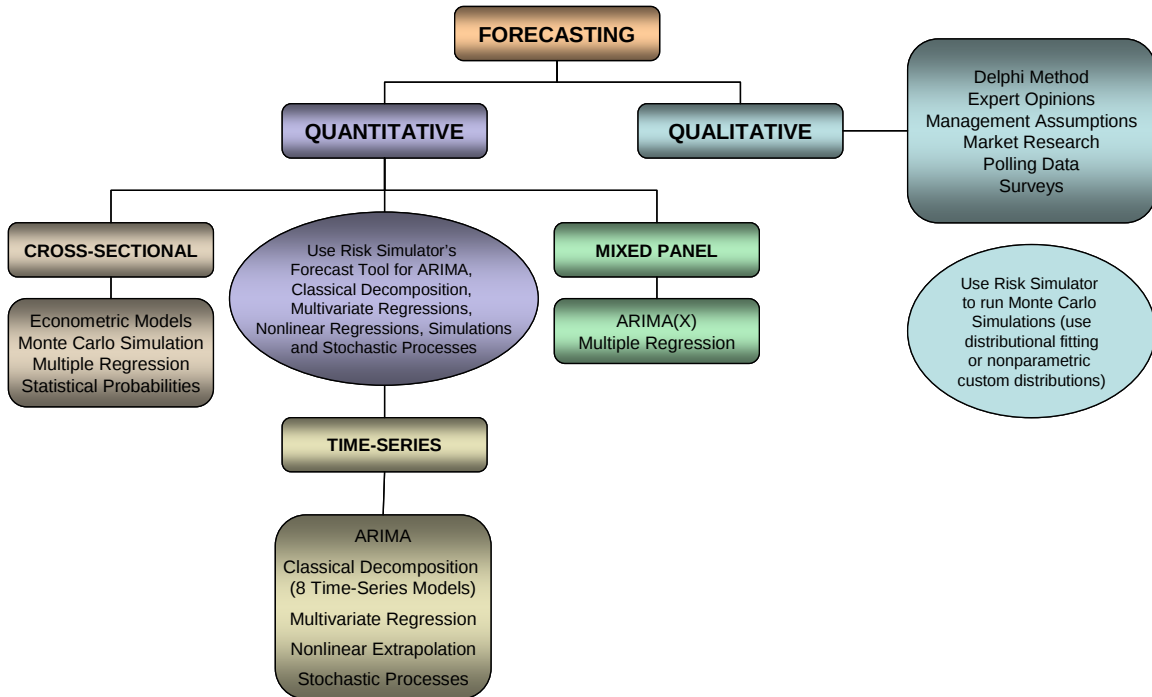


Figure 3.1 – 예측 과정

가

예측법은 대부분 양적, 질적으로 나누는 것이 가능합니다. 질적 예측법은, 확실하지 않고, 또는 존재하지 않는 이력, 같은 시기, 및 비교 데이터가 존재할 때에 사용됩니다. Delphi, 및 전문 의견 (제품 영역의 전문가, 마케팅의 전문가, 혹은 내부 스텝 멤버에 의한 합의 형성), 관리의 가정 (상급 관리자에 의한 목표 성장률의 설정), 마케팅 리서치, 외부 데이터, 혹은 투표, 및 조사 (제 3 자의 정보원, 산업 및 섹터의 지표, 및 적극적 시장 조사등에서 얻은 데이터) 와 같은 복수의 정성적 예측법이 있습니다. 이것의 추정, 단일·포인트의 추정 (평균치의 일치) 가 예측의 집합 설정 (예측의 분포) 의 어느쪽인가의 가능성이 있습니다. 그 후, 리스크 시뮬레이터에 주문 분포를 커스텀 분포를 기입하는 것이 가능, 결과로서 나타난 예측을 시뮬레이션 하는 것이 가능합니다. 이것은, 분포로서 추정된 이것의 데이터 포인트 자체를 사용한 Nonparametric 시뮬레이션입니다.

정략적 예측법에서는, 실재하는 데이터 혹은 예측을 실행하는 것에 필요한 데이터는, 시계열 (다른 년도의 수입, 인플레이션의 비율, 이률, 마켓 쉐어, 고장등의 시간 요소를 가진 값) 、크로스 섹션 (각 학생의 SAT 의 레벨, IQ 와 일주일에 소비하는 알코올 음료의 수가 부여되고, 어떤 특정의 년도에 전국의 대학 2 년생의 평균의 성적에 의한 시간에서 독립한 값) 、또는 혼합 패널 (시계열과 패널 (보조적) 데이터의 혼합、예를 들어、주어진 마케팅의 비용과 마켓 쉐어 하에 차기 10 년간의 매상의 예측、이것은 、매상이 시계열 이라고 하더라도、마케팅의 비용 및 마케팅 쉐어가 예측을 모델화할 때의 지원을 하기위하여 존재하는 것들의 의미합니다) 으로 분할 가능합니다。

데이터

리스크 시뮬레이터의 소프트웨어는 여러가지 예측법을 포함하고 있습니다。

1. ARIMA (자기 회귀 화분 평균 이동)
2. 자동 ARIMA
3. 자동 계량법
4. 기본적인 계량법
5. 큐빅·스플라인
6. 가
- 7.
8. GARCH (일반화 자기 회귀 조건 추가 분산 불균일 모델)
9. J 곡선
10. S 곡선
11. 마르코프 체인
12. 조합 퍼지 로직 예측
13. 최우법
14. 비선형 외압법
15. 회귀 분석
16. 확률 예측법
17. 시계열 분석
18. 트렌트 라인

각 예측법의 분석의 상세는、이 사용자의 매뉴얼의 범위에서 벗어나 있습니다만、자세히는、리스크 시뮬레이터의 소프트웨어의 개발자인 조나단 문 박사에 의한 리스크 의 모델화 : 몬테카를로 시뮬레이션、리얼 옵션 분석, 확률 예측법과 포트폴리오의 최적화의 적용(Wiley Finance 2006)을 참조 부탁드립니다。그것과 또、하기의 무엇보다도 일반적인 어프로치가 표시되어 있습니다。다른 모든 예측법의 어프로치가 꽤 간단히 리스크 시뮬레이터에서 적용하는 것이 가능합니다。

일반적으로 예측을 작성하는 것은 복수의 스텝을 행할 필요가 있습니다.

- Excel 을 기동하여 기입하거나, 실재하는 이력적 데이터를 열어 주십시오.
- 데이터를 선택하고, 시뮬레이션 (Simulation) 、 예측 (Forecasting) 을 선택하여 주십시오.
- 적절한 범위 (ARIMA, 다변수 회귀, 비선형 외삽법, 확률 예측, 시계열 분석)을 선택하고, 적절한 입력을 기입하여 주십시오.

Figure 3.2 는, 예측 툴과 여러 가지 방법을 표시하고 있습니다.

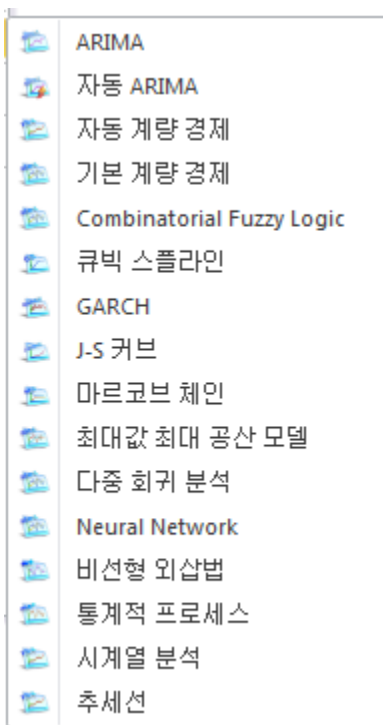


Figure 3.2 – 리스크 시뮬레이터의 예측 과정

다음에 각 방법의 개략적 검토와 다음의 각 방법의 개략적 검토와 소프트웨어를 사용하는데 있어서 복수의 간결한 예증이 부여됩니다. 예증 파일은, 스타트 메뉴에서 여는 것이 가능합니다. *스타트 (Start) | 리얼 옵션즈 밸류에이션 (Real Options Valuation) | 리스크 시뮬레이터(Risk Simulator) | 예증(Examples)*을 선택하거나, *리스크 시뮬레이터(Risk Simulator) | 예증 모델 (Example Models)*을 선택하여 주십시오.

이론:

Figure 3.3 는, 계절성과 경향에 의하여 분리된 여덟개의 무엇보다도 일반적인 시계열 모델을 표시하고 있습니다. 예를 들어, 데이터 변수가 경향, 및 계절성을 가진 경우, 단일 이동 평균 모델, 또는 단일 지수 평활법 모델에서 충분할 것입니다. 단, 계절성이 존재하거나 확실한 경향이 존재하지 않는 경우는, 계절성의 가법, 또는 계절성의 상승적 모델의 적용이 더욱 적절하다는 것입니다.

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

Figure 3.3 – 여덟개의 무엇보다도 일반적인 시계열법

순서:

- Excel 을 기동하고, 필요한 경우는 이력적 데이터를 열어 주시기 바랍니다 (하기의 예증에서는 예증 필터에 있는 시계열 예측의 파일을 사용하고 있습니다).
- 이력적 데이터를 선택하여 주십시오 (데이터는 일렬로 표시되지 않으면 안됩니다)
- **리스크 시뮬레이터(Risk Simulator) /예측 (Forecasting) /시계열 분석(Time-Series Analysis)**을 선택하여 주십시오.
- 적용할 모델을 선택하고, 관련 된 가정을 기입하여, OK 를 클릭하여 주십시오.

시계열 분석

기준치 시계열, 트렌드, 계절성 계수에 대한 알파, 베타, 감마 파라미터를 찾아내기 위한 최적화 프로세스를 모델에 적용하고 예측에서 이러한 것들을 보여줍니다. 즉 저를 "Backcast"에 적용 할 방법론은 최적의 피팅을 찾아내고 예측 오류를 최소화하기 위한 모델의 best fitting 파라미터를 찾을 후 Exit하여, 이력적 데이터에 기초한 미래의 "forecast"를 진행합니다. 동일한 기준치의 성장, 트렌드, 앞으로 진행되는 계절성에 대하여 가정하는 것이 이것의 코스입니다. 만약 구조적 시프트가 존재할 때 (예 : 글로벌화, 합병, 파생 등과 같은) 그렇게 말하지 않더라도, 기준치 예측은 계산 되어 줄 수 있고 요구 되는 수정사항이 예측을 만들어 낼 수 있습니다.

Procedure

이 모델을 실행하기 위하여 간단히:

1. 이력적 데이터를 선택하십시오 (cells E25:E44).

2. Simulation | Forecasting | Time-Series Analysis를 선택하십시오.

3. Auto Model Selection, Forecast 4 Periods and Seasonality 4 Periods를 선택하십시오.

기존의 시뮬레이션 프로파일 이 있을 경우에는 Create Simulation Assumptions만 선택 하실 수 있습니다. (그렇지 않을 경우, Simulation을 클릭하십시오.)

신규 시뮬레이션 프로파일, 각 스럽에 대한 시계열 예측을 진행하고, Create Simulation을 클릭하십시오.

가장 박스).



Year	Quarter	Period	Sales
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	\$677.00
2007	3	7	\$1,008.60
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	\$1,280.40
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 - 시계열 분석

결과의 해석에 있어서:

Figure 3.5 는, 예측 툴을 통하여 생성된 샘플 결과를 표시하고 있습니다. 사용된 모델은, Holt-Winters 의 상승적 모델입니다. Figure 3.5 에서는, 적합 모델과 예측 차트는, 경향과 계절성은 Holt-Winters 의 상승적 모델에 의하여 잘 선택되는 것을 표시하고 있는 것에 주목하여 주십시오. 시계열 분석의 레포트는, 중요한 최적화된 알파, 베타, 감마의 파라미터, 에러의 측정, 적합한 데이터, 예측값과 적합한 예측의 그래프를 부여하여 주고 있습니다.

파라미터는 단순히 기본적으로 표시되고 있습니다. 알파는, 종료 시간에 변환하는 기준 레벨의 기억의 효과를 취득하고, 베타는, 경향의 강함을 측정하는 경향 파라미터로써, 감마는, 이력적 데이터의 계절성의 강함을 측정합니다. 분석은, 이력적 데이터를 이러한 세 개의 요소로 분해하고, 미래의 예측을 실행하기 위하여 다시 구성합니다. 적합한 데이터는, 구성의 모델을 사용하고 적합한 데이터로서 이력적 데이터를 표시하고, 과거에 대하여 어느 정도 예측이 가까운가 (이 기술은 백 캐스팅이라고 불리워 집니다)을 표시하고 있습니다. 예측의 값은, 단일 포인트의 추정인가 가정인가 (시뮬레이션 프로파일 존재하고, 가정의 자동 생성이 선택되지 않는 경우) 어느쪽인가 입니다. 그래프는, 이것의 이력적, 적합한

예측의 값을 표시하고 있습니다. 차트는 강력한 커뮤니케이션으로서, 예측 모델은 어느 정도 좋은가를 구분하는 시각 툴이기도 합니다.

메모:

이 시계열 분석의 모듈에서는, Figure 3.3 에 표시된 여덟개의 시계열 모델이 포함되어 있습니다. 경향과 계절서의 기준에 기초한 특정의 모델을 선택하고 실행하는 것이 가능, 또는, 자동적으로 여덟개의 모델과 반복적으로 조정하고, 파라미터를 최적화 하여, 데이터에 무엇보다도 적절한 모델을 추출하는 자동 모델 셀렉션을 선택하는 것도 가능합니다. 한편, 여덟개의 모델 중에서 한 개의 모델을 선택하는 경우, 최적화의 체크 박스의 선택을 끝내는 것이 가능, 독자의 알파, 베타, 그리고 감마 파라미터를 입력하는 것이 가능하게 됩니다. 이것의 파라미터의 기술적 상세에서는, 조나단 문 박사의 리스크 모델화 : 몬테카를로 시뮬레이션의 적용, 리얼 옵션스 분석, 예측과 최적화 (Wiley, 2006) 를 참조하십시오. 또한, 자동 모델 셀렉션, 및 계절성 모델의 어느쪽인가를 선택한 경우, 적절한 계절성 기간을 입력할 필요가 있습니다. 계절성의 입력은 플러스의 정수가 아니면 안됩니다 (예, 데이터가 4 기의 경우, 년간의 계절, 사이클을 수에 4로 입력하거나, 데이터가 월간의 경우에는 12로 입력하여 주십시오). 다음에 예측하는 기간의 수를 기입하여 주십시오. 이 값은, 플러스의 정수가 아니면 안됩니다. 최고의 실행 시간은 300 초입니다. 보통, 변환의 필요가 아닙니다. 단, 꽤 많은 양의 이력적 데이터에 기초한 예측의 분석이 조금씩 길게 되어, 과정 시간이 실행 시간을 넘는 경우, 과정은 종료되어 버립니다. 또는, 자동적 과정의 생성을 행한 예측을 선택하는 것이 가능합니다. 이것은, 단일 포인터에 대한 예측이 가정됩니다. 최후에, 극 파라미터 옵션은, 알파, 베타와 감마 파라미터를 최적화하는 가능성을 부여하고, 0 과 1 을 포함하는 것이 가능합니다. 일부의 예측의 소프트웨어밖에, 이것의 극단적인 파라미터의 사용을 가능하게 해 주지 않습니다. 리스크 시뮬레이터는, 어느 것을 사용하는가 선택할 가능성을 부여하여 주고 있습니다. 일반적으로는 극단의 파라미터를 사용할 필요는 없습니다.

Holt-Winter's

	RMSE		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		

$\alpha = 0.0001$, $\beta = 1.0000$, $\gamma = 6$

gamma 가 (s) (L) 가 , 가 , 가
 가 2 (MSE) , (RMSE) RMSE
 (RMSE) MSE , RMSE
 (MAPE) , Theil U 가 1.0

1	684.20
2	584.10
3	765.40
4	892.30
5	885.40
6	677.00
7	1006.60
8	1122.10
9	1163.40
10	993.20
11	1312.50
12	1545.30
13	1596.20
14	1260.40
15	1735.20
16	2029.70
17	2107.80
18	1650.30
19	2304.40
20	2639.40
예측21	2713.69
예측22	2114.79
예측23	2900.42
예측24	3293.81
예측25	3346.55
예측26	2580.81
예측27	3506.19
예측28	3947.61

RMSE	71.8132
MSE	5157.1348
MAD	53.4071
MAPE	4.50%
Theil's U	0.3054

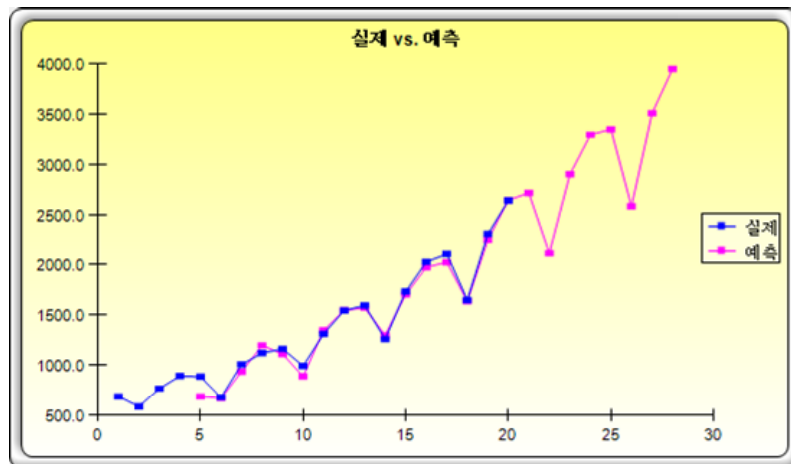


Figure 3.5 – 예증 : Holt-Winters 의 예측 레포트

이론:

여기에서는, 사용자가 회귀 분석의 기초에 의하여 충분한 지식을 가지고 있 충분한 지식을 가지고 있을 필요가 있습니다. 일반적인 2 변수적 선형 회귀의 방식은 $Y = \beta_0 + \beta_1 X + \varepsilon$ 로, β_0 는 절편으로, β_1 는 기울기이며, ε 는 에러를 표시하고 있습니다. 이것은, 두 개의 변수만을 포함한 2 변수 관수로, Y 혹은 종속 변수와, X 혹은 독립 변수가 됩니다. X 는 회귀 변수로도 알려져 있습니다 (시절, 2 변수 회귀는, 독립 변수 X 만이 존재할 때는 단회귀라 합니다). 종속 변수는, 독립 변수에 의존하기 위하여 종속 변수라 불리워지고 있습니다. 예를 들어, 판매의 수입은, 상품의 광고 및 프로모션에서 사용되는 마케팅의 비용에 종속하고 있어, 종속 변수의 판매와 독립 변수의 마케팅의 비용을 작성하고 있습니다. 2 변수 회귀의 예증은, Figure 3.6 의 좌측의 패널에 표시되어 있는 2 차원면에서 데이터 포인트의 설정을 통하여, 최적의 적합선의 삽입입니다. 다른 케이스에서는, 복수, 또는 독립한 X 변수의 n 수의 다변수 회귀를 실행할 수 있습니다. 보통의 회귀 방식은 $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \dots + \beta_n X_n + \varepsilon$ 입니다. 이 케이스는, 최적의 적합선은, $n + 1$ 차원면내에 포함됩니다.

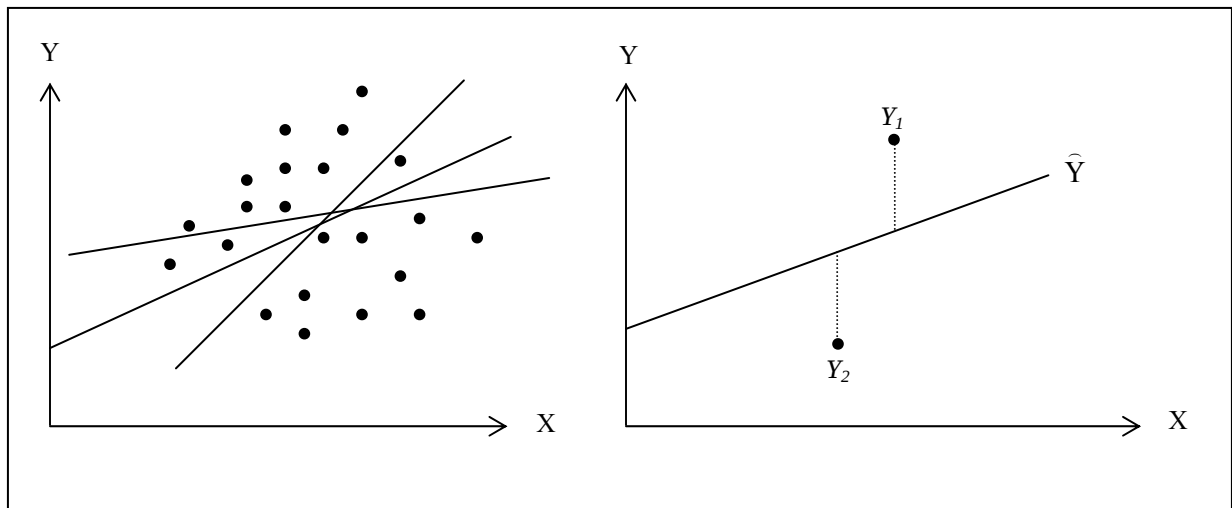


Figure 3.6 – 이변량 회귀

단, Figure 3.6 와 같은 데이터 포인트의 설정을 통하여 선을 적합히 하면, 결과로서 복수의 확률적 선이 표시되고 있습니다. 최적의 적합선은, 종합적인 종속의 에러를 최소화하는 단일적인 특이한 선으로 정해져 있습니다. 이것은, Figure 3.6 의 우측의 패널에서 표시되어 있는 것과 같이, 현재의 데이터 포인트(Y)와, 예기되어 있는 선 ($\hat{Y}$)의 사이의 절대 거리의 덧셈입니다. 에러를 최소화하고 최적의 적합 선을 내기 위하여, 무엇보다도 상급적인 어프로치가 필요합니다. 이것은 회귀 분석입니다. 따라서 회귀 분석은, 전제적으로 에러를 최소화하는 산식에 의하여 특이한 적합선을 냅니다.

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

여기에서는, 하나의 특이한 선만이 편차평방을 최소화합니다. 에러 (현재의 데이터 포인트와 예기된 선 사이의 종축의 거리) 는, 플러스 에러를 마이너스 에러가 캔슬된 것을 피하기 위하여 평방 처리됩니다. 기울기와 절편에 관하여 최소화의 문제 해결에는, 1 차 관수를 계산하여, 제로에 가깝게 되도록 계산하지 않으면 안됩니다:

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

이 결과는, 2 변수 회귀의 최소 이승식을 가져옵니다:

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

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

다변수 회귀에도, 복수의 독립 변수를 고려하여 유의의 확장을 행하는 것이 가능하고, 더불어, $Y_i = \beta_1 + \beta_2 X_{2,i} + \beta_3 X_{3,i} + \varepsilon_i$ 로, 그 결과 추정 경사는 다음과 같이 계산됩니다:

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

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

다변수 회귀를 실행하는 것에 있어서, 결과의 설정과 해석에는, 주의가 필요합니다. 예를 들어, 계량 경제학적인 모델화 시에는 뛰어난 이해가 (예, 구조의 파탄, 중공선성, 불균일 분산, 자기 상관, 특정한 테스트, 비선형등과 같은 검출)적절한 모델의 구축에 필요합니다. 다변수 회귀의 분석과 논의 및 어떻게 하여 이것의 함정을 검출하는가의 상세에는, 조나단 문 박사의 리스크 모델화: 몬테카를로·시뮬레이션의 적용, 리얼 옵션 분석, 예측과 최적화 (Wiley, 2006) 를 참조하십시오.

순서:

- Excel 을 기동하고, 필요한 경우 이력 데이터를 열어 주십시오 (하기의 표에서는, 예증 파일의 **복수의 회귀**를 사용하여 주십시오).
- 편집된 데이터가 열에 있는 것을 확인하여 주십시오. 변수 명칭을 포함한 모든 데이터 범위를 선택하고, **리스크 시뮬레이터(Risk Simulator) / 예측(Forecasting) / 복수의 회귀 (Multiple Regression)**를 선택하여 주십시오.

- 중속 변수를 선택하고, 중요한 옵션 (타임랙, 스텝 와이즈 회귀, 비선형 회귀 등)을 선택하고, OK 를 클릭하여 주십시오.

결과의 해석:

Figure 3.8 는 , 샘플의 다변수 회귀 결과의 레포트를 표시하고 있습니다. 레포트는 모든 회귀 결과, 분산 분석 결과, 적합 차트와 가설 검정의 결과가 포함되어 완전한 정보를 표시하고 있습니다. 이것의 결과의 해석 기법의 상세는, 이 사용자 매뉴얼의 최후에 있습니다. 회귀 레포트의 해석과 같이 다변수 회귀의 논의와 분석의 상세는, 조나단 문 박사의 리스크 모델하: 몬테카를로, 시뮬레이션의 적용, 리얼 옵션즈 분석, 예측과 최적화 (Wiley 2006) 를 참조하여 주십시오.

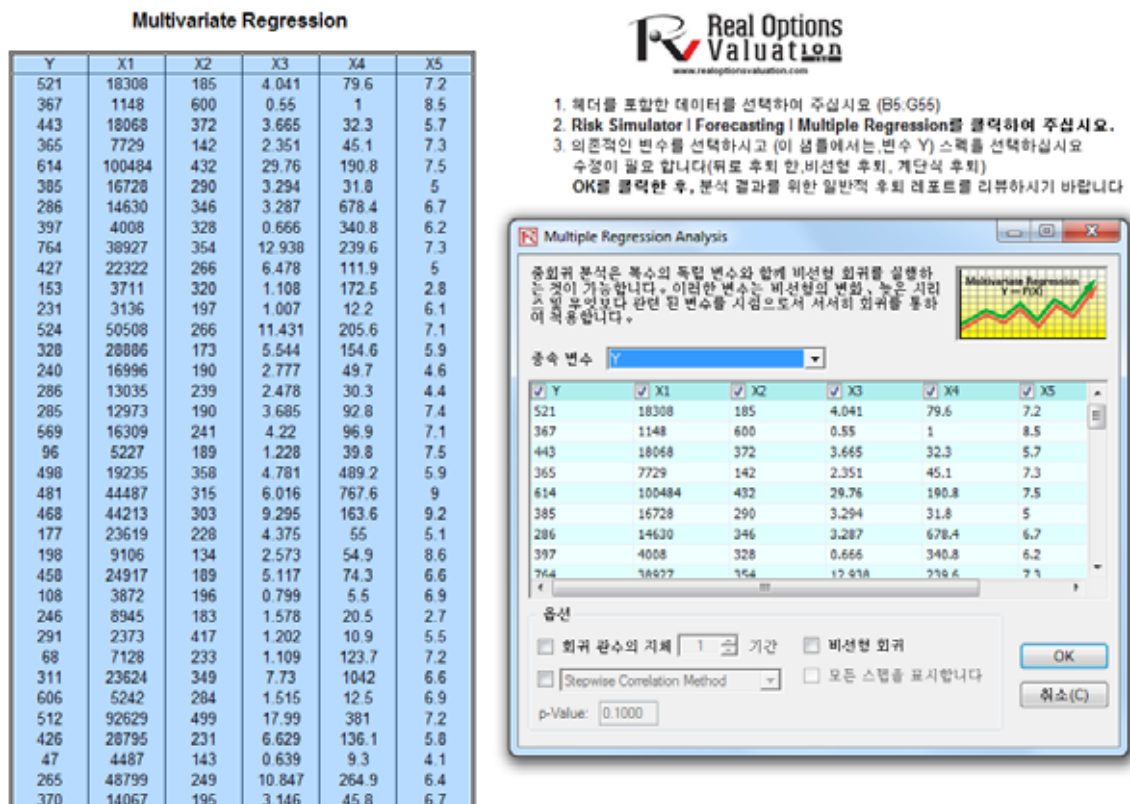


Figure 3.7 – 다변량 회귀의 실행

R-Squared ()	0.3272
R ()	0.2508
(SEy)	0.5720
	149.6720
	50

(R - Squared) 、 、 、 (R - Squared) 、 가 가
 (Multiple R) 、 (Y) 、 (Y)
 (R - Square) 、
 (SEy) 、 、

	X1	X2	X3	X4	X5	
t-통계	57.9555	-0.0035	0.4644	25.2377	-0.0086	16.5579
p-값	108.7901	0.0035	0.2535	14.1172	0.1016	14.7996
5%보다 낮음	0.5327	-1.0066	1.8316	1.7877	-0.0843	1.1188
95%보다 높음	0.5969	0.3197	0.0738	0.0807	0.9332	0.2693
	-161.2966	-0.0106	-0.0466	-3.2137	-0.2132	-13.2687
	277.2076	0.0036	0.9753	53.6891	0.1961	46.3845

5	t- (99%	44	)	2.6923
44	t- (95%	44	)	2.0154
49	t- (90%	44	)	1.6802

$b_1X_1 + b_2X_2 + \dots + b_nX_n$. 가 . $t-$. $Y = b_0 +$

t-통계는, 실질적인 상관의 평균값 = 0의 귀무 가설(Ho), 실제의 평균값이 0이 아니고, 대립 가설(Ha)이 설정된, 가설 검정으로 사용됩니다. T-검정이 실행되어, 계산된 t-통계는, 비판적인 값으로서 관련 된 잔차의 자유도로 비교됩니다. 다른 회귀를 대상으로, 각 상관이 통계적으로 유의한 경우, t-검정 그 자체와, t-검정의 계산은 매우 중요한 것이 됩니다. 이것은, t-검정은 회귀 관수, 또는 독립 변수가 회귀에 남아 있어야 하는가, 혹은 Drop되어야 할 것인가를 통계적으로 확인합니다.

t-	(df)	t- 가	p-
90%, 95% 99%			
0.01, 0.05 0.10	99%, 95% 99%		
p-	90%	0.10	

F-	p-	가	F-	(99%	5	44)	)	3.4651
479388.49	95877.70	4.28	0.0029	F-	(95%	5	44)	2.4270
985675.19	22401.71			F-	(90%	5	44)	1.9828
1465063.68								

(ANOVA) 、 F - 、 T - 、 F -
 가 가 p - 、 가 가 (Ho) 、 F - 、
 가 (Ha) 、 0.01, 0.05, 0.10 가 F - 、 p - F

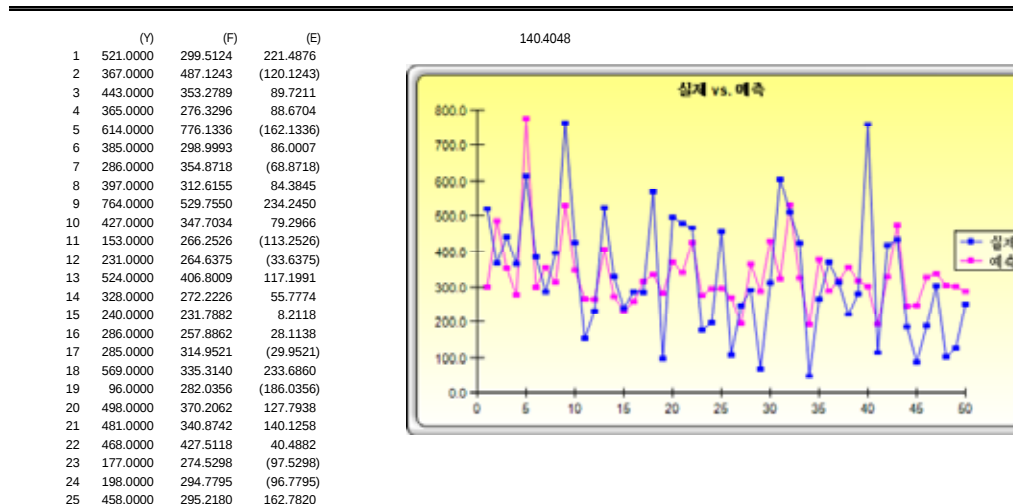


Figure 3.8 – 다변량 회귀의 결과

이론:

확률 과정은, 시간이 지나가는 것에 따라 작성된 결과의 시리지의 수학적 공식밖에 없습니다. 이 결과는 성질로서 결정론이 아닙니다. 이것은, 매년 상승하는 X 퍼센트 가격 및 이 X의 요소에서 나오는 Y 퍼센트에 의하여 상승하는 수입과 같이 명확하고 심플한 룰에 따른 방식, 또는 과정은 아닙니다. 확률 과정이란, 비결정론인 것에서 불리워, 확률 과정의 공식에 수치를 이어주는 것이 가능하고, 매회 다른 결과를 얻는 것이 가능합니다. 예를 들어, 주식가격의 루트는 성질로서 확률적이 되기 위하여, 확률적 주식 가격의 루트를 예측하는 것이 가능합니다. 단, 시간이 지나는 것에 따라 진화하는 가격은, 이것의 가격을 생성하는 과정으로 둘러싸이게 됩니다. 과정은 결과와 다르게, 사전에 고정되고, 정해지게 되어 있습니다. 따라서, 확률 시뮬레이션에 의하여, 가격의 복수의 경로를 작성하고, 이것의 시뮬레이션의 통계적인 샘플링을 취득하고, 시계열의 생성에 사용되는 확률 과정의 파라미터와 성질을소요하는 경우에 실제의 가격이 얻는 잠재적인 경로에 관하여 추측하는 것이 가능합니다. 세계의 기본적인 확률 과정이 리스크 시뮬레이터의 예측 툴에 포함되어 있어, 먼저, 기하 그라운 운동 혹은 램덤 워크는, 단순히 폭넓은 응용에 의하여 무엇보다도 일반적으로 넓게 사용되고 있습니다. 다른 두 개의 확률 과정은, 평균 회귀 과정과 점프 확산 과정입니다.

확률 과정 시뮬레이션으로 흥미 깊은 곳은, 이력적 데이터가 반드시 필요하지 않다는 것입니다. 이것은, 모델은 이력 데이터의 집합에 적합할 필요가 없는 것입니다. 단, 이력 데이터의 예측 변동률과 기대 리턴을 하는 것만이, 외부 데이터로 비교하여 이것을 추정하거나, 이것의 값의 가정을 작성하는 것이라는 것입니다. 어떻게 하여 각 입력이 계산 (예, 평균 회귀의 비율, 점프 확률, 예측 변동률 등) 되는가 등의 상세는, 조나단 문 박사의 리스크 모델화: 몬테카를로 시뮬레이션의 적용, 리얼 옵션즈 분석, 예측과 최적화, 제 2판 (Wiley 2006)을 참조하여 주십시오.

순서:

- **리스크 시뮬레이터(Risk Simulator) / 예측 (Forecasting) / 확률 과정 (Stochastic Processes)**를 선택하는 것으로서 모듈을 기동시켜 주십시오.
- 희망하는 과정을 선택하여, 필요로 하는 입력을 기입하여 주십시오. 수 회 차트를 갱신하도록 클릭하여, 과정이 희망대로 돌아가는 것을 확인하고, OK를 클릭하여 주십시오. (Figure 3.9)

결과의 해석:

Figure 3.10는, 샘플의 확률 과정을 표시하고 있습니다. 차트는, 반쪽의 샘플 설정을 표시하고 있어, 레포트는, 확률 과정의 기초를 기술하고 있습니다. 또는, 각 시간의

시간을 위하여 예측값 (평균값과 표준 편차) 가 부여되고 있습니다. 이것의 값의 사용에 의하여, 분석에 있어서 어느 시간 주기가 가장 중요한 것인가를 결정하는 것이 가능하고, 정규 분포를 사용한 이러한 평균과 표준 편차의 값에 기초한 가정의 설정이 가능합니다. 이것의 가정은 차후에, 주문 모델에서 시뮬레이션 하는 것이 가능합니다.

확률 과정은, 확률법에서 생성된 통속사상, 또는 연속 퍼드의 것을 말합니다. 이것은, 랜덤 사상은 장기 시간에 발생하는 것이 있습니다만, 특성의 통계적인, 그리고 확률적인 물에 의하여 정하여 집니다. 주요 확률 과정은, 랜덤 워크, 혹은 브라운 운동, 평균 회귀와 점프 확산이 포함되어 있습니다. 이러한 과정은, 봤을 때 랜덤한 경향을 보이지만, 확률법에 제한되어 있는 변수의 변수를 예측하는 것이 가능합니다. 리스크 시뮬레이터의 확률 가장 모델에서 각 과정을 작성, 시뮬레이션 할 수 있습니다. 이 프로세스는 주가, 물가, 오일 가격, 전력가격, 원자재 가격 등의 가격 변동과 같이 다수의 변수가 있는 시계열 변수의 예측에 사용하는 것이 적절합니다.

이 모델을 실행하기 위해서는, 간단히:

1. Simulation | Forecasting | Stochastic Processes를 선택합니다.
2. 관련 입력값의 세트를 입력하거나 테스트 케이스의 기존의 입력값을 사용하여 주십시오.
3. 시뮬레이션 하기 위하여 관련 프로세스를 선택 하십시오.
4. 단일 경로의 계산될 계산을 보기 위하여 업데이트 차트를 클릭 하거나, 프로세스 생성을 위하여 OK를 클릭하십시오.



Figure 3.9 – 확률 과정의 예측

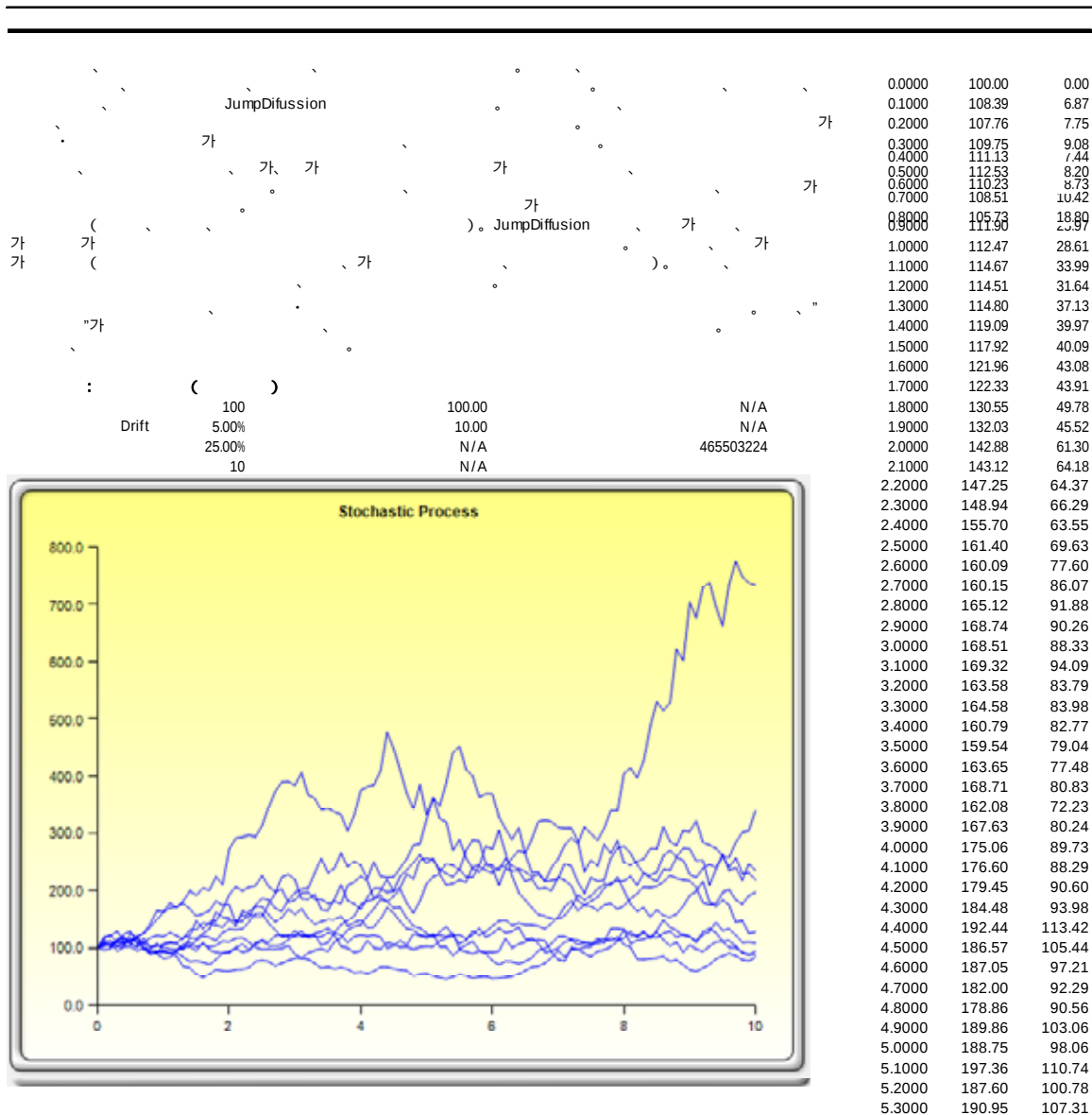


Figure 3.10 - 확률 예측법의 결과

이론:

외압법은, 미래의 특정의 기간에 예측된 이력적 경향의 사용에 의하여 통계적인 예측을 작성합니다. 이것은, 시계열만으로 사용하고 있습니다. 크로스 섹션, 또는 혼합 패널의 데이터 (클로스 섹션 데이터를 활요한 시계열)에 의하여, 다변수 회귀가 적절합니다. 이 방법은, 커다란 변화가 기대되지 않는 경우에 최적입니다. 이것은, 원인 요소가 일정하게 유지 되는 것을 기대합니다, 또는 상황의 원인 요소가 명확히 알려 지지 않을 경우에 사용하는 것이 최적입니다. 또는, 과정에의 개인적인 바이어스에의 도입을 방지하는 것에 도움이 됩니다.

외압법은 꽤 신뢰할 수 있어서, 비교적 심플하고, 경제적입니다. 단, 근황, 그리고 이력적 경향이 지속한다는 가정의 외압법은, 프로젝트의 기간내에 불연속이 발생한 경우, 커다란 예측 에러를 나타냅니다. 즉, 시계열의 순수한 외압법은, 예측 된 시리즈의 이력값에 모든 필요한 정보가 포함되어 있다고 가정하고 있습니다. 과거에 일어나 것이 미래에 일어난 다는 것을 예언하는 것에 최적이라는 가정을 한다면, 외압법이 최적입니다. 따라서, 모든 필요한 것은, 많은 짧은 예측이라는 시기에, 매우 사용하기 쉽고, 적절한 어프로치입니다.

이 기법은, 모든 임의의 x 값을 통하여 매끄러운 비선형 곡선을 투입하고, 이력적 데이터 집합에 미래의 x 값을 외압하는 것에 의하여 $f(x)$ 의 관수를 추정합니다. 이 기법은, 다항식의 관수 형태, 또는 로지스틱 공식의 형태 (두 개의 다항식 비율) 의 어느쪽인가를 사용합니다. 일반적으로, 데이터가 잘 돌아가는 것을 얻기 위하여, 다항식의 관수 형태로 충분합니다만, 로지스틱 관수 형태는, 때에 따라 정도를 보여 주는 것이 있습니다 (특히 극단수의 경우. 예, 분모가 제로에 접근하는 관수) .

순서:

- Excel 을 기동하고, 필요한 경우는 이력적 데이터를 열어 주십시오 (다음에 표시되고 있는 일러스트는 예증 폴더의 **비선형 외압법**의 파일을 사용하고 있습니다).
- 시계열 데이터를 선택하고, **리스크 시뮬레이터 (Risk Simulator) /예측 (Forecasting) 비선형 외압법 (Nonlinear Extrapolation)** 을 선택하여 주십시오.
- 외압법의 타입을 선택하고(자동 선택, 다항식 관수, 및 유리 관수), 희망하는 예측 기간의 수를 기입하고 (Figure 3.11), OK 를 클릭하여 주십시오.

결과의 해석:

Figure 3.12 에 표시되고 있는 결과는, 외압된 예측의 값, 에러의 측정법과 외압법의 결과의 그래프를 표시하고 있습니다. 에러의 측정법은, 예측의 타당성을 확인하기 위하여 사용 할 수 있고, 예측의 성질을 비교할 때에 무엇보다도 중요하게 되어, 시계열 분석에 대하여 외압법의 정도를 비교할 때에 적절합니다.

메모:

이력적 데이터가 매끄럽게 몇 개의 비선형 패턴과 곡선에 따를 경우, 외압법은 시계열분석보다 최적입니다. 단, 데이터 패턴이 계절성의 싸이클과 경향을 따를 경우는 시계열 분석쪽이 좋은 결과를 표시합니다.

외삽법은, 미래의 특정한 기간의 상황을 이력적 경향을 사용하는 것으로, 통계적인 예상을 포함합니다. 이것은, 시계열 예측법으로밖에 사용되지 않습니다. 외삽법은 신뢰성 있고, 비교적 심플하며, 저렴합니다. 그러나, 최근의 이력적 경향을 계속하기 위한 외삽법은 프로젝트 시간 종료안에서 종단이 이루어질 경우 거대한 예측 에러를 불러 옵니다. projected time period.

Historical Sales Revenues Polynomial Growth Rates

Year	Month	Period	Sales
2010	1	1	\$1.00
2010	2	2	\$6.73
2010	3	3	\$20.52
2010	4	4	\$45.25
2010	5	5	\$83.59
2010	6	6	\$138.01
2010	7	7	\$210.87
2010	8	8	\$304.44
2010	9	9	\$420.89
2010	10	10	\$562.34
2010	11	11	\$730.85
2010	12	12	\$928.43



1. 이력적 데이터의 입력과 데이터 범위를 선택하십시오 (E13:E24)
- 2 Risk Simulator | Forecasting | Nonlinear Extrapolation을 클릭하십시오.
3. Function type을 선택하고 외삽법 기간이 필요합니다. 클릭 OK

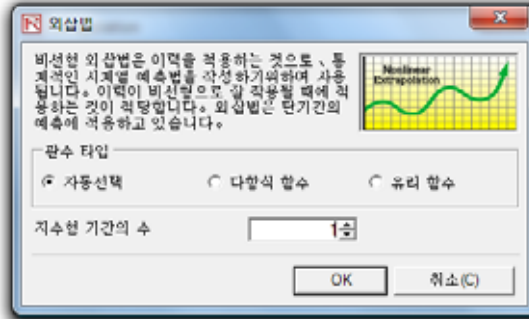


Figure 3.11 – 비선형 외삽법의 실행

비선형 외삽법

통계의 전략

외삽법은, 미래의 특정한 기간의 상황을 이력적 경향으로 사용하는 것에 의하여, 통계적 예상을 포함합니다. 이것은, 시계열 예측법만으로 사용되지 않습니다. 횡단, 및 혼합 패널 데이터 (시계열과 횡단 데이터) 에는, 종회귀가 최적입니다. 이 방법은 커다란 변화가 예상되지 않을 때에 적절합니다. 이것은, 원인 요인은 남아있는 정수에예기되어 있으나, 원인 요인의 상태가 확실히 이해되지 않을 때를 표시합니다. 그것은 아직 프로세스내의 개인적인 선입관예의 도입의 실망 등을 지원합니다. 외삽법은 확실히 신뢰되고, 비교적 사용하기 쉽고, 저렴합니다. 단, 최신, 및 이력적 경향은 연속이라고 가정하는 외삽법은, 반영된 시간내에서 불연속이 발생하면, 장기의 예측 에러가 표시됩니다. 즉, 시계열의 순수한 외삽법은, 알 필요가 있는 모든 데이터는, 예측 된 시리즈의 이력적 가치에 포함되어 있다고 가정할 수 있습니다. 과거의 이력적 동향이 미래의 동향에 유망한 예언이라고 가정하는 경우, 외삽법은 대단한 권위입니다. 이것은, 모든 데이터를 단 기간 예측하여야 할 때에 적절합니다.

이 방법은, 모든 x의 가치에 의하여 원만하게 비선형 커브의 삽입에 의하여 모든 임의의 x가치를 위하여 f(x)함수를 추정하고, 이 원만한 커브를 사용하고, 역사적 데이터 세트를 넣고 미래의 x의 가치를 외삽법으로 추정합니다. 방법은 다항식 함수형식이 이성적인 함수 형식(두 개의 다항식의 비율)을 필요로 합니다. 보통, 다항식 함수 형식은, 좋은 동향을 가진 데이터를 위하여 충분합니다만, 이성적 함수 형식에 의하여 보다 정확한 때가 있습니다(특히 극함수와, 예를 들어, 제로에 가까운 분모와 함수들).

기간	현재	적합한 예측
1	1.00	
2	6.73	1.00
3	20.52	-1.42
4	45.25	99.82
5	83.59	55.92
6	138.01	136.71
7	210.87	211.96
8	304.44	304.43
9	420.89	420.89
10	562.34	562.34
11	730.85	730.85
12	928.43	928.43
예측13		1157.03
예측14		1418.57
예측15		1714.95

에러의 측정	
RMSE	19.6799
MSE	387.2974
MAD	10.2095
MAPE	31.56%
Theil의U	1.1210

기능 타입: 이성적

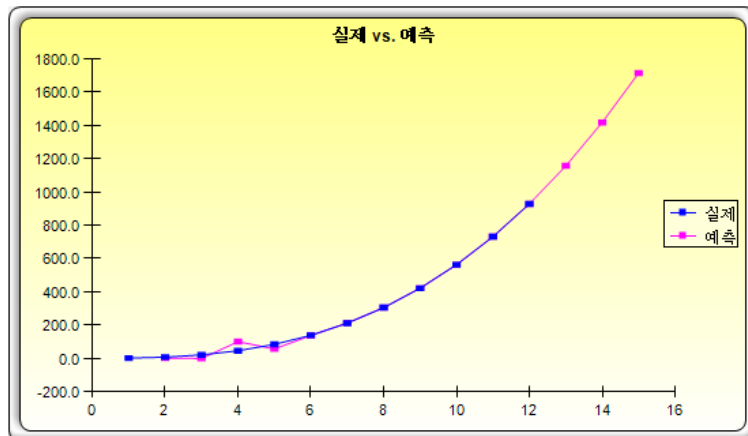


Figure 3.12 – 비선형 외삽법의 결과

Box-Jenkins ARIMA

이론:

하나의 매우 고도의 강력한 시계열 예측 툴인 ARIMA, 또는 자기 회귀 화분 평균 분석입니다. ARIMA 의 예측은 세 개의 분리된 툴을 알기 쉬운 모델을 모았습니다. 최초의 툴 세그먼트는, 자기 회귀, 및 “AR”로, 조건이 없는 예측 모델에 남는 타임랙의 값의 수에 대응합니다. 본질적으로는, 모델은 예측 모델에의 현재의 데이터의 이력적 변동을 채취하고, 이 변동 및 나머지 값을 사용하여 더 우수한 예측 모델을 만들어 냅니다. 두 번째 툴 세크먼트는 화분으로서 “I”로 표시하고 있습니다. 이 화분은, 예측 된 시계열을 미분하는 회수에 해당합니다. 이 요소는, 데이터에 존재하는 모든 비선형 성장의 비율을 카운트 합니다. 세 개의 툴 세그먼트는, 이동 평균, 및 “MA”로 표시되고 있어, 본질적으로 시간차를 가진 예측 에러의 이동 평균을 표시합니다. 이 시간차를 갖은 예측의 에러를 포함한 것으로 모델은, 본질적으로 이것의 예측 에러, 또는 차이에서 배우는, 이동 평균의 계산을 통하여 수정을 하였습니다. ARIMA 모델은, Box-Jenkins 방법에 따라서, 모델 구성에서 얻은 각 스텝을 Random Noise 만이 남을때까지 행하여 집니다. 즉, ARIMA 의 모델화는, 예측의 생성에 상관을 사용합니다. ARIMA 는, 구상된 데이터에서 표시되지 않는 모델 패턴도 사용할 수 있습니다. 또는, ARIMA 모델은, 외생변수와 혼합하는 것이 가능하나, 외생변수가, 추가의 예측 기간을 예측하는 것에 충분한 데이터 포인트를 가지고 있는 것을 확인하여 주십시오. 최후에, 모델의 복잡함을 위하여, 이 모듈은 성장 실행에 있는 것을 이해하여 주십시오.

ARIMA 모델은 일반적으로 시계열 분석과 다변수 회귀보다 고도의 이유가 복수 있습니다. 일반적으로 시계열 분석과 다변수 회귀의 이해는, 남아 있는 에러는 이것 자신의 시간차를 가진 값의 상관입니다. 이 연속적인 상관은, 방해와 다른 방행와 상관하지 않는 다는 회귀 이론의 표준적인 가정에 위반합니다. 연속 상관에 관련한 주요 문제에는 다음과 같은 점이 생각됩니다.

- 회귀 분석과 기본적인 시계열 분석은 다른 선형의 추정값 사이에서는 별로 유효하지 않습니다. 단, 에러의 잔여가 현재의 에러의 잔여를 예측하는 것에 도움이 됨으로, ARIMA 를 사용하고 종속 변수의 보다 좋은 예측에 이 정보를 활용하는 것이 가능합니다.
- 회귀와 시계열 방식을 사용하고 계산된 표준 에러는 바르지 않고, 일반적으로 적당히 서술되어, 만약, 회귀 예측값으로서 시간차가 있는 종속 변수가 존재하는 경우, 회귀의 추정에 바이어스가 있어, 수미일관하지 않으나, ARIMA 를 사용하는 것에 의하여 수정하는 것이 가능합니다.

자기 회귀 화분 이동 평균, 및 ARIMA(p,d,q) 모델은, 시계열 데이터에서 연속 상관의 모델화를 위한 세 개의 요소를 사용하는 AR 모델의 연장입니다. 최초의 컴포넌트는 자기 회귀(AR)입니다. AR(p)모델은, 공식에서 시계열의 p 의 시간차를 사용합니다.

AR(p)모델의 방식은 $y_t = a_1 y_{t-1} + \dots + a_p y_{t-p} + e_t$ 입니다. 두 번째의 컴포넌트는 화분 (d)입니다. 각 화분 $j \geq 0$ 이라는 y_{t-j} ; 는, 시계열을 미분하는 것에 대응합니다. I(1)는, 데이터를 한번 미분하는 것을 의미합니다. I(d)는, d 회의 데이터의 미분을 의미하고 있습니다. 세 번째의 컴포넌트는 이동 평균 (MA)입니다. MA(q) 모델은, 예측을 향상시키기 위하여 예측 에러의 시간차 q 를 사용합니다. MA(q)모델의 방식은 $y_t = e_t + b_1 e_{t-1} + \dots + b_q e_{t-q}$ です. 최후에, ARIMA(p,q) 모델은 결합 형태 $y_t = a_1 y_{t-1} + \dots + a_p y_{t-p} + e_t + b_1 e_{t-1} + \dots + b_q e_{t-q}$ 을 가지고 있습니다.

순서:

- Excel 을 기동하고, 데이터 기입하거나 실재하는 워크 시트와 이력적 데이터를 예측하는데 열어 주십시오 일러스트에서는, 다음의 예증에서 사용하는 파일 **시계열 ARIMA** 을 표시하고 있습니다).
- 시계열 데이터를 선택하고, **리스크 시뮬레이터 (Risk Simulator) / 예측(Forecasting) / ARIMA** 을 선택하여 주십시오.
- 중요한 P, D, 와 Q 이 파라미터 (플러스의 정수만)을 입력하고, 희망하는 예측 주기의 수를 기입하고, OK 를 클릭하여 주십시오.

결과의 해석:

ARIMA 모델의 결과의 해석에는, 다변수의 회귀 분석의 지정과 거의 같습니다.

(ARIMA 모델과 다변수의 회귀 분석의 해석 기법의 상세는 조다만 문 박사의 리스크 모델화, 제 2 판을 참조 하여 주십시오) 。 단, Figure 3.14 에서 표시된 것과 같이 ARIMA 분석에 특정한 결과의 여러가지 부가 설정이 있습니다. 처음에, ARIMA 모델의 선택, 및 동일 증명에서 잘 사용되는 아카이케 정보 기준 (AIC)과 슈왈츠 기준 (SC)의 부가입니다.

이것, AIC 와 SC 는, 특정의 p, d 와 q 의 파라미터를 가진 특정의 모델이 최적인 경우 통계가 어떠한가를 정하기 위하여 사용됩니다. SC 는, AIC 보다 부가계수를 위한 커다란 패널티를 부여하고 있으나, 일반적으로는, AIC 와 SC 의 값이 낮은 모델이 선택됩니다. 최후에, 자기 상관 (AC)이라 불리는 결과의 부가 설정과 부분 자기 상관 (PAC)통계는, ARIMA 레포트에서 부여 됩니다.

예를 들어, 자기 상관 AC(1)이 제로가 아닌 경우는, 시리즈는 일차 연속적 상관을 의미합니다. 단, AC 가 기하학적 시간차의 증가로써 쇠퇴하는 경우는, 시리즈는 더 낮은 자기 회귀 과정에 따르고 있는 것을 표시합니다. 만약, 소수의 시간차의 후에 AC 가 제로로 하강하는 경우는, 시리즈는 낮은 이동 평균 과정에 따르는 것을 표시합니다.

한편, PAC 은, 개입하는 시간차에서 상관을 제거한 후에 k 주기 떨어진 값의 상관의 측정을 합니다. 상관의 패턴은 k 보다 적은 차수의 자기 회귀에 의하여 채워지는 경우는, 시간차 k 에 있는 부분 상관은 제로에 가깝습니다. Ljung-Box Q-통계와 시간차 k 에 대하여 이것의 p-값은 이미 부여되어 있어, 검정된 귀무 가설은, k 의 차수까지 자기 상관을 볼 수 없습니다. 자기 상관의 구상에서 그려진 점선은, 크게 2 표준 오차의 범위내입니다. 만약 자기 상관이 이것들의 한계내에

들어가 있지 않은 경우 및 5%의 유의 수준에서 제로와는 어떠한 유의한 상관도 인정되지 않습니다.

바른 ARIMA 모델 발견에서는, 연습과 경험이 필요합니다. 이것의 AC, PAC, SC, 와 AIC 는, 바른 모델의 사양을 검출하는 매우 유용한 진단 툴입니다.

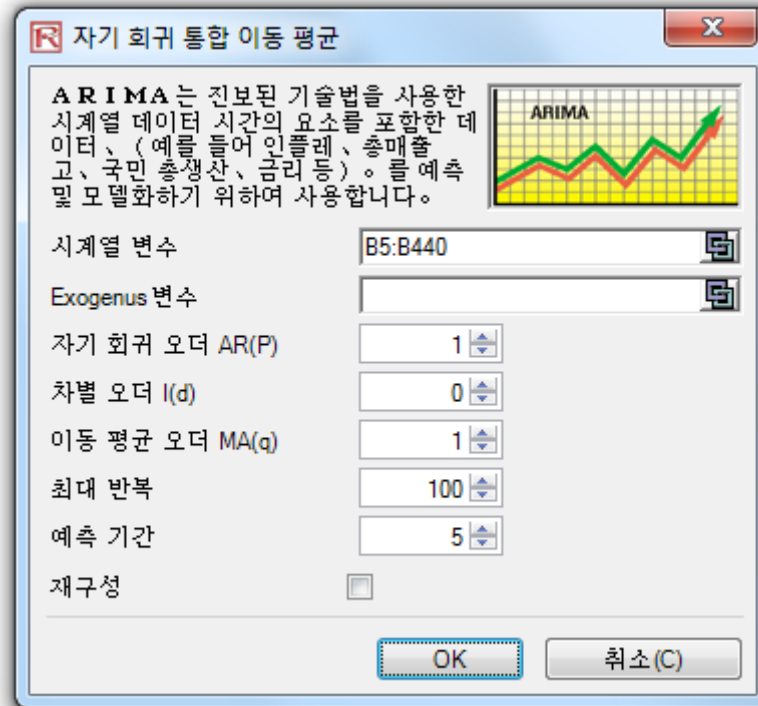


Figure 3.13 Jenkins box ARIMA 의 예측 툴

R -	(	)	0.9999	(AIC)	4.6213
R -			0.9999	Schwarz (SC)	4.6632
R(		)	1.0000	L	-1005.13
	(SEy)		297.52	Durbin-Watson (DW)	1.8588
			435		5

가	(99%	X df)	2.5873
	(95%	X df)	1.9655
t-	(90%	X df)	1.6484

(ANOVA) , F , t - , F , 가 , 가 , f - , 가 , 가 , 가 (Ha) , 가 , p - , 가 (Ho) , p - , 0.01, 0.05 , 0.10 , F- , 가

	(Y)	(F)	(E)
2	139.4000	139.6056	(0.2056)
3	139.7000	140.0069	(0.3069)
4	139.7000	140.2586	(0.5586)
5	140.7000	140.1343	0.5657
6	141.2000	141.6948	(0.4948)
7	141.7000	141.6741	0.0259
8	141.9000	142.4339	(0.5339)
9	141.0000	142.3587	(1.3587)
10	140.5000	141.0466	(0.5466)
11	140.4000	140.9447	(0.5447)
12	140.0000	140.8451	(0.8451)
13	140.0000	140.2946	(0.2946)
14	139.9000	140.5663	(0.6663)
15	139.8000	140.2823	(0.4823)
16	139.6000	140.2726	(0.6726)
17	139.6000	139.9775	(0.3775)
18	139.6000	140.1232	(0.5231)
19	140.2000	140.0513	0.1487
20	141.3000	140.9862	0.3138
21	141.2000	142.1738	(0.9738)
22	140.9000	141.4377	(0.5377)
23	140.9000	141.3513	(0.4513)
24	140.7000	141.3939	(0.6939)
25	141.1000	141.0731	0.0270
26	141.6000	141.8311	(0.2311)
27	141.9000	142.2065	(0.3065)
28	142.1000	142.4709	(0.3709)
29	142.7000	142.6402	0.0598
30	142.9000	143.4561	(0.5561)
31	142.9000	143.3532	(0.4532)
32	143.5000	143.4040	0.0960
33	143.8000	144.2784	(0.4784)
34	144.1000	144.2966	(0.1966)
35	144.8000	144.7374	0.0626
36	145.2000	145.5692	(0.3692)
37	145.2000	145.7582	(0.5582)
38	145.7000	145.6649	0.0351
39	146.0000	146.4605	(0.4605)
40	146.4000	146.5176	(0.1176)
41	146.8000	147.0891	(0.2891)
42	146.6000	147.4066	(0.8066)
43	146.5000	146.9501	(0.4501)
44	146.6000	147.0255	(0.4255)
45	146.3000	147.1382	(0.8382)
46	146.7000	146.6328	0.0672
47	147.3000	147.4819	(0.1819)

Auto ARIMA (Box-Jenkins ARIMA)

이론:

이 틀은, ARIMA 모듈과 같은 분석을 부여하고 있으나, 차이로써는, Auto · ARIMA 모듈은, 모델의 사양의 복수의 구성 조합의 자동적인 검증으로 몇 개의 전형적인 ARIMA 를 자동화하여, 최적으로 적합한 모델을 보여 줍니다. Auto·ARIMA 의 실행은 보통의 ARIMA 예측에 유의하여 있습니다. 차이로서는, P, D, Q 의 입력은, 이미 요구되지 않고, 이것의 입력의 다른 조합은 자동적으로 실행되고 비교됩니다.

순서:

- ❏ Excel 을 기동하고, 데이터를 입력하거나 실재하는 워크 시트와 예측하기 위한 이력 데이터를 열어 주십시오 (Figure 3.14 로 표시되는 일러스트는, 리스크 시뮬레이터의 **예측** 메뉴에 있는 **고도의 예측 모델**의 파일을 사용하고 있습니다).
- ❏ **자기 ARIMA** 워크 시트를 **리스크 시뮬레이터 (Risk Simulator) / 예측(Forecasting) / 자기·ARIMA (AUTO-ARIMA)** 를 선택하는 것으로 열어 주십시오.
- ❏ 링크 아이콘을 클릭하고, 존재하는 시계열 데이터를 관련시켜, 희망하는 예측 주기를 입력하여 **OK**를 클릭하여 주십시오.

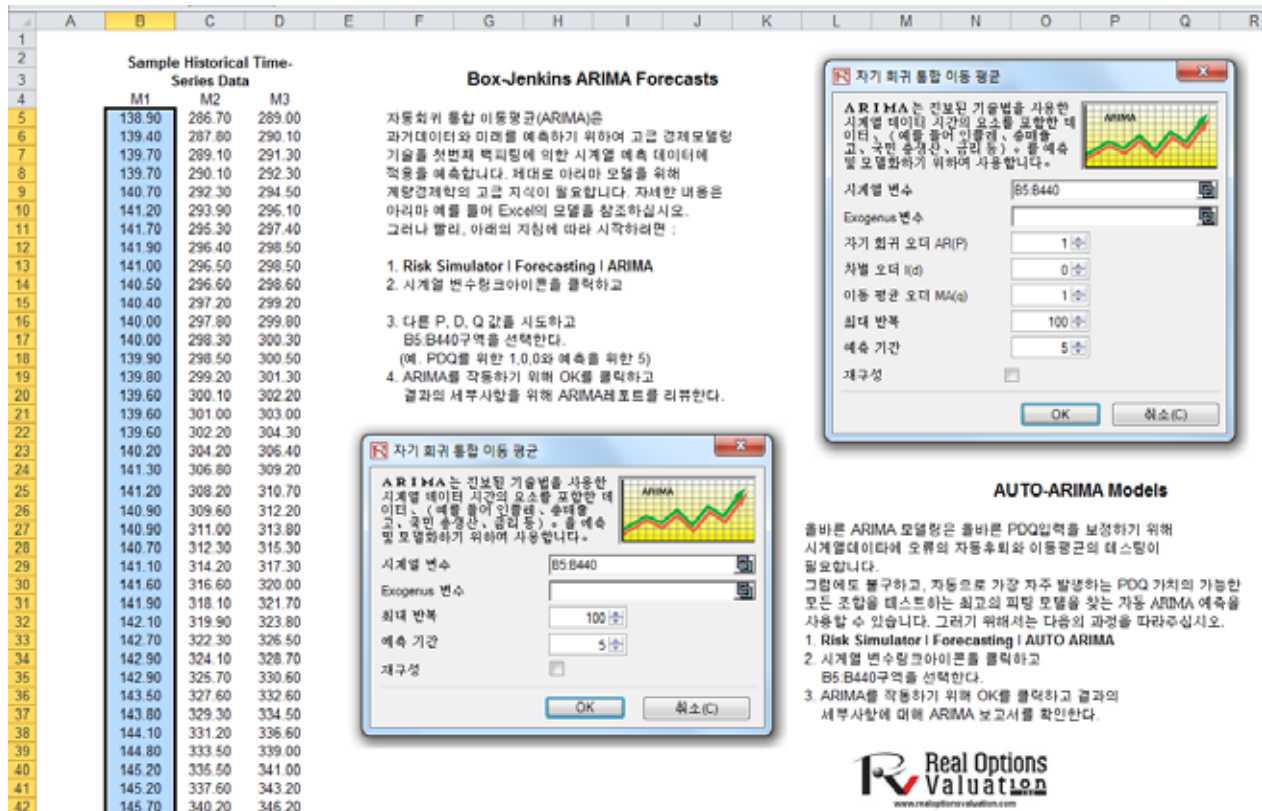


Figure 3.14 AUTO ARIMA 모듈

이론:

계량 경제학은, 어떤 비즈니스·경제 변수의 변화의 모델화 혹은 예측을 위한 비즈니스 분석, 모델화와 예측 기법의 한 분야의 것을 의미합니다. 기본적으로는 계량 경제학의 모델을 실행하는데, 종속, 및 독립 변수는 회귀 실행 이전에 변환하는 것 이외에는, 보통의 회귀 분석과 같습니다. 생성된 레포트는 중회귀의 장에서 표시되는 것과 같이, 해석 기법은, 이전 기술 된 것과 같습니다.

순서:

- Excel 을 스타트하고, 데이터를 기입하거나, 실재하는 워크 시트로 예측하기 위한 이력 데이터를 열어 주십시오 (Figure 3.15 로 표시되는 일러스트는, 리스크 시뮬레이터의 **예증** 메뉴의 **고도의 예측 모델**의 파일을 사용하고 있습니다).
- 기본적인 계량 경제** 의 워크 시트를 데이터로 선택하여, **리스크 시뮬레이터 (Risk Simulator) | 예측(Forecasting) | 기본적인 계량 경제(Basic Econometrics)**를 선택하여 주십시오.
- 희망하는 종속, 및 독립 변수를 (Figure 3.15 의 예증을 참조하여 주십시오) 기입하고, **OK** 를 클릭하고, 모델과 레포트를 실행하여 주십시오. 또는, **결과의 표시**를 클릭하고, 모델내에서 무언가의 변환을 행하는 경우는 레포트의 작성 이전에 결과를 보는 것이 가능합니다.



Figure 3.15 기본적인 계량 경제학의 모델

J-S

이론:

J-곡선, 또는 지수 성장 곡선은, 차주기의 성장이 현재의 주기 레벨에 의존하여, 증가가 지수 관수적에 있는 곡선의 것입니다. 이것은, 시간이 흐르는 것에 따라서, 한 개의 주기에서 다음의 주기까지의 값은, 꽤 증가한다는 것을 의미하고 있습니다. 이 모델은, 일반적으로 시간이 변화에 있어서 생물의 성장과 화학 반응의 예측에 사용됩니다.

순서:

- Excel 을 기동하고, **리스크 시뮬레이터 (Risk Simulator) / 예측(Forecasting) / JS 곡선(JS Curves)** 을 선택하여 주십시오.
- J, 또는 S 곡선의 타입을 선택하고, 필요한 입력 가정을 기입하여 주십시오 (Figures 3.16 와 3.17 을 예증으로서 참조항 주십시오) 。 모델과 레포트를 실행하는 것의 **OK**를 클릭하여 주십시오.

수확에서는, 지수관수적에 성장하는 양은 성장률이 현재의 사이즈에 항상 비례되고 있습니다. 그러한 성장은 지수 법칙에 계속된다고 불러지고 있습니다. 이것은 온갖 지수 관수의 성장달은, 양이 보다 커지면, 그것만 빠르게 성장하는 것을 의미합니다. 그러나 종속 변수의 사이즈와 성장률의 사이의 관계는 엄밀한 법칙에 의하여 지배되는 것을 표시합니다. 또한 무엇보다도 간단한 표현으로는 직접 비율을 표시합니다. 지수 성장의 일반 원칙에서는, 수가 보다 크면 흡수율, 보다 빠른 성장이 보여집니다. 어떤 지수 관수적으로 크게 되는 수치도, 같은 시간내에서, 일정한 비율만 크게 되는 다른 숫자보다 크게 될 가능성이 있습니다. 이 예측 방법은 J의 자에 현상이 유의하게 있는 것에서 J-곡선으로도 불러워지고 있습니다. 이 성장곡선의 최고 레벨은 없습니다. 다른 성장 곡선은 S곡선 및 Markov Chains이 포함되어 있습니다.

J-커브 예측을 만들기 위하여, 아래의 지시를 따라 주십시오:

1. 클릭 Risk Simulator | Forecasting | JS Curves
2. Exponential J Curve를 선택하고 원하는 값을 입력하십시오
(예: Starting Value를 100, Growth Rate 5 퍼센트, 종료End 기간100)
3. 예측을 위해서는 OK를 실행하고 forecast 레포트의 리본을 위하여 수 분 기다립니다.

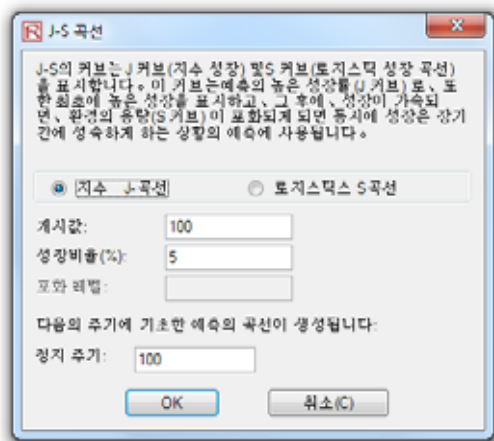


Figure 3.16 J-곡선의 예측

S-곡선, 또는 로지스틱 성장 곡선은, 지수 성장률을 가지고, J 곡선과 같이 개시합니다. 시간이 흘러감에 따라, 환경이 포화 (예, 시장 포화, 경쟁, 혼잡 등)하고, 성장은 늦어져, 예측값은 결국, 포화, 또는 최고 레벨에서 정지하게 됩니다. 이 모델은 일반적으로, 신상품의 시장에서 도입까지 성숙, 저하의 사이클에

관한 마켓 쉐어 혹은 판매 성장의 예측에 사용됩니다. Figure 3.17 는, 샘플 S-곡선을 표시하고 있습니다.

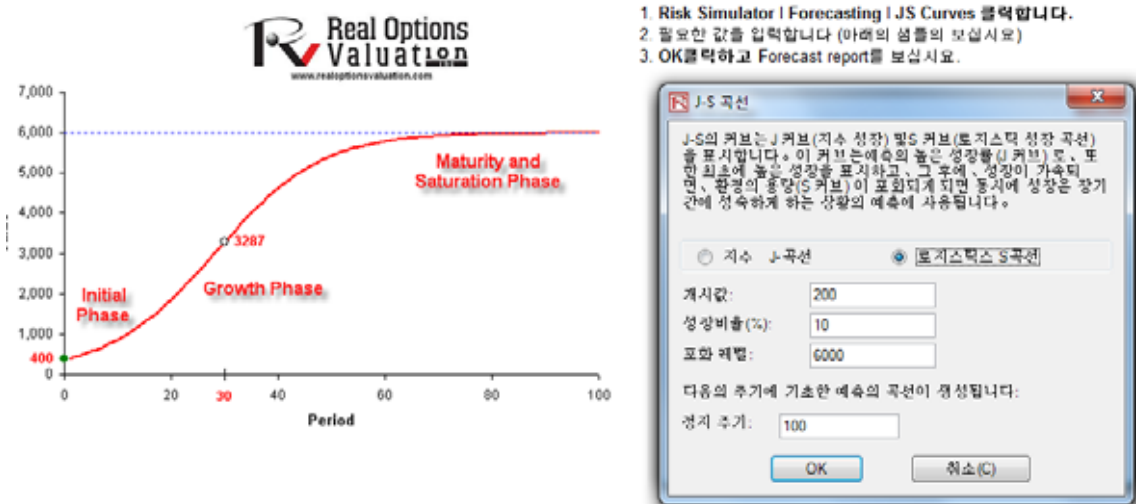


Figure 3.17 S-곡선의 예측

GARCH

이론:

생성된 자기 회귀 조건부 불균일 분산(GARCH) 모델은, 그리고 유가 증권의 이력 예측변동률 레벨의 예측 (예, 주가 가치, 물가, 석유 가격 등) 을 위하여 사용됩니다. 데이터의 집합은, 미가공의 가격 레벨의 시계열이 아니면 않됩니다. GARCH 는, 먼저 가격을 상대적 리턴으로 변환시키고, 이력 데이터가 평균 회귀 변동률의 기간 구조에 적합하게 하기 위하여 내부의 최적화를 실행합니다. 이 과정이 행하여 질 때에, 예측 변동율은 성질로서 불균일성 분산 (어떤 계량 경제학의 성질에 적합한 시간적 변화) 라고 가정합니다. GARCH 모델의 이론적인 상세는, 이 사용자 매뉴얼의 범위외에 있습니다. GARCH 모델의 상세는, 조나단 문 박사의 고도 분석 모델” (Wiley 2008)을 참조 하십시오.

순서:

- Excel 을 기동하고, 예증 파일의 **고도의 예측 모델**을 선택하고, **GARCH** 워크시트를 선택하고, **리스크 시뮬레이터(Risk Simulator) | 예측(Forecasting) | GARCH**를 선택하여 주십시오.
- 링크 아이콘을 클릭하고, 데이터의 배치를 선택하여, 필요한 입력 가정을 기입하고 (Figure 3.18 를 참조 하십시오), **OK** 를 클릭하고, 모델과 레포트를 실행하여 주십시오.

메모: 일반적인 예측 변동률의 예측에 필요한 상황은, $P = 1$, $Q = 1$, 주기 = 년가의 주기수 (월간 데이터는 12, 주가 데이터에는 52, 일차 데이터에는, 252 또는, 365), 기초 = 1 의 최소값과 주기값의 상한과 예측 주기 = 희망하는 년차화된 예측 변동률의 예측 수입니다.



Generalized Autoregressive Conditional Heteroskedasticity (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

GARCH 모델을 실행하려면, 해당 시계열 데이터를 입력한 다음 리스크시뮬레이터 | 리스크시뮬레이터 | 예측 | GARCH를 클릭하고 데이터 위치 링크 아이콘에 클릭하고, 과거 데이터 영역을 선택합니다. (예, C8 : C2428) 필수 입력 (예, P 1, Q 1, 일일 거래 주기성 252, 예측 자료 1, 예측 기간 10)을 입력하고 확인을 누릅니다. 생성된 예측 보고서를 리뷰합니다.

GARCH가 일반화된 자기회귀 조건 포함의 불균일 분산 모델은, 가격 자신을 사용하여 금융 상품의 변동률의 예측에 사용됩니다. GARCH (P,Q) 모델은, 평균(뉴스) 및 변동공식의 다른 플러스 P 및 정수 Q의 지체된 파라미터를 가능하게 합니다. 플러스 데이터 가치만이 GARCH의 변동률 예측에 사용하는 것이 가능하다는 것에 주의하여 주십시오. 주기단 1년 사이의 주기를 숫자로 표시합니다 (예를 들어, 12, 252, 매일 데이터는 365). 이것은 변동률을 연간적으로 표시하기 위하여, 혹은 주기적인 변동률을 1의 주기로서 갖기 위한입니다. 기반은, 추정된 기초 주기입니다 (이것은 예측의 기반으로 미래의 변동률을 추정하는데 몇 개의 과거의 주기가 필요한가를 표시하고 있습니다. 대부분 1에서 12 사이의 값입니다.) 변동의 목표는, 만약 변동률의 예측을 기입한 긴 주기의 평균종료 시간을 늘리기 위한 때에 표시합니다. 미가공 가격 데이터를 연대순(과거에서 현재의 값을 다수의 열로 단일 컬럼으로 표시)에 정리하는 것을 확인하여 주십시오.

데이터 위치:

GARCH (P,Q) 모델을 일반화:

P: Q: 주기: 기본: 예측 기간:

☐ 분산 타겟을 적용하십시오

☒ GARCH
 ☐ GARCH-M
 ☐ TGARCH
☐ TGARCH-M
 ☐ EGARCH
 ☐ EGARCH-T
☐ GJR GARCH
 ☐ GJR TGARCH
 ☐ 모든 모델 실행

OK 취소

Figure 3.18 GARCH 예측 변동률의 예측

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

(Markov Chains)

이론 :

Markov chain 는, 미래 상태의 확률이, 이전의 상태에 의존하여, 그와 같은 관계가 함께 링크된 체인의 형상을 하고 있고, 장기간적인 정상 상태의 레벨을 검수할 때에 존재합니다. 이 어프로치는 일반적으로 2 개사의 경쟁 마켓 웨어 예측을 하는 때에 사용됩니다. 필요한 입력은, 처음의 스토아 (초기 상태) 에 간 소비자가 다음에 같은 스토아에 가는 확률에 대하여, 경쟁 스토아에 갈 천이 확률입니다.

순서:

- Excel 을 기동하고, *R 리스크 시뮬레이터 (Risk Simulator) | 예측(Forecasting) 마르코프 체인(Markov Chain)*을 선택하여 주십시오.
- 필요한 입력 가정을 기입하고 (예증에서는 Figure 3.19 을 참조하여 주십시오), **OK**를 클릭하고 모델과 레포트를 실행하여 주십시오.

Markov과정은, 연속적인 기한에 시스템의 배수상으로, 및 반복된 시행의 진화를 조사하기 위하여 유용합니다. 특정시의 시스템 상태는 미지로, 특정의 상태가 있는 것의 확률을 아는 것에 주목합니다. 예를 들어 특정 기계 및 장치가 다음의 주기에 작동하여 계속되거나, 혹은, 소비자 구입 상품 A가 다음의 주기 이내에 상품 A를 구입을 계속할까, 혹은 경쟁 브랜드 B에 전환할 것인가 등의 확률을 계산하는데, MarkovChains가 사용됩니다.

Markov 프로세스를 만들기 위하여, 아래의 지시를 따라 주십시오:

1. 클릭 **Risk Simulator | Forecasting | Markov Chain**
2. Relevant state probabilities 을 입력하고 (예: 90과 80퍼센트) **OK**를 클릭하십시오
3. Forecast report 리뷰하십시오

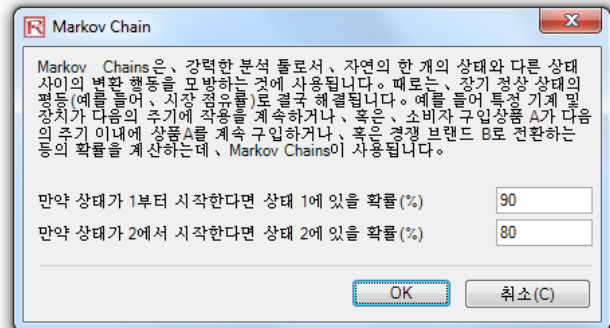


Figure 3.19 Markov Chains 천이 체제

(MLE: Logit, Tobit, Probit)

이론 :

유한 종속 변수는, 2 값 반응 (0, 1), 절단되거나, 정연된, 혹은, 절단된 데이터와 같은 스코프와 범위가 유한인 종속 변수에 데이터가 포함되어 있는경우의 시추에이션을 기술합니다. 예를 들어, 종속 변수의 데이터 세트를 부여한 경우(예, 연령, 수입, 신용카드, 혹은 저당 대부권자의 교육 레벨), 최고 우도 추정 (MLE)를 사용하여 책임 불이행의 확률을 모델화하는 것이 가능합니다. 반응, 혹은 종속 변수 Y 가 2 값인 경우, 1 과 0 의 2 업만으로서의 확률적인 결과밖에 할 수 없습니다(예, Y 는, 특정의 조건의 존재/부재, 이전의 론의 책무 불이행/부·책임 불이행, 어떤 디바이스의 성공/실패, 개관의 네/아니오의 대답 등). 혹은, 종속 변수 회귀 X 의 벡터를 소지하고 있는 것으로 결과 Y 를 영향하는 것으로 정의되어 있습니다. 무엇보다도 일반적인 최소 2 승 회귀 어프로치는, 회귀 에러가 불균형 분산성으로 있는 비정상·정상임으로, 무효로, 측정되는 확률 측정은, 1 이상, 0 미만의 무의미한 값이 결과로서 나옵니다. MLE 분석은, 이것의 문제를 종속 변수가 유한인 경우에 로그의 우도 기능을 최대화하기 위한 반복 루틴의 최적화 사용을 조작합니다.

Logit, 혹은 로지스틱 회귀는, 데이터를 로지스틱 곡선에 적합한 것으로 이벤트의 발생 확률을 예측하기 위하여 사용됩니다. 이것은, 2 항식 회귀를 위하여 사용되는 일반화된 선형 모델로써, 여러가지 회귀 분석의 형식과 같이, 수적, 혹은 분류적으로 어느 쪽이 아니면 안되는 여러가지 예측 변수를 이용하고 있습니다. 2 값의 다중 로지스틱 분석에서 적용된 MLE 는, 특정의 그룹에 속하는 예측된 성공 확률을 정의하기 위한 종속 변수를 모델화 하기 위하여 사용됩니다. Logit 모델을 위하여 측정된 계수는, Logarithmic 확률의 비율로써, 확률이라고 불리는 솔직한 해석은 아닙니다. 처음으로 간단한 계산을 필요로 하고, 어프로치는 간단합니다.

명확히는, Logit 모델은, 측정된 $Y = \text{LN}[P_i/(1-P_i)]$, 혹은, 역으로 말하면, $P_i = \text{EXP}(\text{측정된 } Y)/(1+\text{EXP}(\text{측정된 } Y))$ 로써 정의된, 계수 β_i 는, 로그 확률의 비율임으로, 안티 로그, 혹은, $\text{EXP}(\beta_i)$ 를 취하는 것으로, $P_i/(1-P_i)$ 의 확률의 비율을 얻는 것이 가능합니다. 이것은, β_i 의 Unit 을 증가하면 로그의 확률의 비율도 이 값에서 증가하는 것을 표시합니다. 최후에는, 확률에서의 변환의 비율은, $dP/dX = \beta_i P_i(1-P_i)$ 입니다. 표준 에러는, 예측된 계수가 어느 정도 정확한 것인가를 측정하고, t 통계는, 이러한 표준 에러에의 예측된 각 계수의 비율로써, 측정된 각 파라미터의 중요성의 일반적인 회귀 가설검정을 위하여 사용됩니다. 특정의 그룹에 속하기 위하여 성공 확률을 측정하는 것은(예, 일년간 피운 담배의 개수를 준, 흡연자가 폐기관의 병기를 발생하는가 어떤가를 예측), 측정된 Y 값을 MLE 계수를 사용하여, 계산합니다. 예를 들어, $Y=1.1+0.005(\text{담배의 개수})$ 가 모델인 경우, 일년간에 100 팩을 소비하는 사람은, $1.1+0.005(100)=1.6$ 가 측정된 Y 와 기술됩니다. 다음에, 확률의 비율의 안티 로그를

계산합니다. 이것은, $\text{EXP}(\text{측정된 } Y)/[1+\text{EXP}(\text{측정된 } Y)]=\text{EXP}(1.6)/(1+\text{EXP}(1.6))=0.8320$ 처럼 표시됩니다. 따라서, 그 예증에 상당하는 인물은, 83.20%의 확률에서 폐기관의 병기를 발생하는 가능성을 가지고 있습니다.

Probit 모델은(Nomit 모델로도 알려져 있습니다), Probit 회귀로 불리워지는 어프로치의 최대 우도 측정을 사용한 Probit 기능을 이용하는 2 값 대응 모델을 대신하는 인기 있는 서술입니다. Probit 과 Logistic 모델은, 매우 상이한 예측을 생성하고, 로지스틱 회귀에서 측정된 파라미터는, 적절한 어프로치 모델보다도 1.6에서 1.8 회 이상의 경향을 표시합니다. Probit 과 Logit 의 어느쪽을 사용하는가는 상화에 따라, 무엇보다 명확한 구별은, 로지스틱 분포는, 극도한 값을 계산하기 위하여, 높은 우도(Fat · Tail)을 가지고 있습니다. 예를 들어, 모델화의 대상이 주인의 판단에 의하여, 이 반응 변수는, 2 값으로(집의 구입, 혹은 집을 구입하지 않음), 수입, 연령등의 종속 변수 X_i 의 시리즈상에 종속되어, $li=\beta_0+\beta_1X_1+...+\beta_nX_n$ 로 표시하고, li 의 값이 클수록, 집주인의 확률이 높아 집니다. 각 가족에, 중요한 1^* 의 슬로프를 존재하게 해, 만약, 그 값을 넘은 경우는, 집이 구입되고, 역으로, 넘지 않은 경우에는, 집이 구입되지 않습니다. 결과 확률(P)은, 정규적으로 분배되고 있다고 정의되고 있어, $P_i=\text{CDF}(li)$ 와 같은 표준 정규 축적 분포 관수(CDF)가 사용됩니다. 따라서, 측정된 계수를 회귀 모델과 측정된 Y 값을 모두 같이 사용하고, 표준 정규 분포를 적용하는 것이 가능합니다(Excel 의 NORMSDIST 기능, 혹은, 리스크 시뮬레이터의 분배적인 분포 툴을 정규 분포를 선택하여, 평균값을 0 과 표준 편차를 1 로 설정하여 이용하여 주십시오). 최후에는, Probit, 혹은, 확률적인 Unit 의 측정을 얻는 것은, $li+5$ 로 설정합니다(이것은, 예를 들어 확률이 $P_i < 0.5$ 라도, 측정된 결과가 마이너스의 숫자가 되면, 정규 분포는, 0 의 평균값의 주변에서는 대조적입니다).

도빗 모델(단절 도빗)은, 계량 경제학과 바이오 메타릭스의 모델화 기법으로, 비·부의 종속 변수 Y_i 와 한개, 혹은 그 이상의 독립 변수 X_i 의 사이의 관계를 기술하기 위하여 사용됩니다. 도빗 모델은, 계량 경제학 모델로서, 종속 변수는 잘려져 있습니다. 종속 변수가 단절되어 있는 것은, 0 보다 낮은값은 관측되지 않아서 입니다. 도빗 모델은, 숨겨진 관측할 수 없는 변수 Y^* 가 존재하는 것을 정의합니다. 이 변수는, β_i 계수의 벡터를 통해 X_i 변수상에서 선형적으로 종속되어 있어서, 상호관계를 정의합니다. 또한, 정규적으로 분배된 에러의 개요 U_i 가 있어서, 그 상호관계상의 랜덤한 영향을 취득합니다. 관측할 수 있는 변수 Y_i 는, 그려져 있는 변수에 상당하기 위하여 정의되어 있어서, 예를 들어, 그려져 있는 변수가 0 을 상회한다고 해도, Y_i 는, 0 으로써 정의되어 있습니다. 이것은, $Y_i=Y^*, \text{if: } Y^*>0$ 로, $Y_i=0, \text{if: } Y^*=0$ 이기 때문입니다. X_i 상에서 관측할 수 있는 Y_i 의 상호 관계의 파라미터 β_i 는, 통상의 최소 제곱 회귀의 사용에 의하여 측정되며, 회귀 측정의 결과는, 모순이 많음으로, 이율의 하향 방향의 바이어스 된 슬로프 계수와 상회의 바이어스된 방해가 표시되어 있습니다. MLE 만이, 도빗 모델에 조화되어, 도빗 모델에서는, 시그마로 불리는 보조적인 통계가 존재하고, 표준적인 최소 제곱 회귀에서의 측정의 표준 에러에 상당하고, 측정된 계수는, 회귀 분석과 같이

상용됩니다.

순서:

- Excel 을 기동하고, 예증 파일의 **고도의 예측 모델, MLE** 의 워크 시트를 열고, 헤더를 포함한 데이터 가정을 선정하고, **리스크 시뮬레이터 (Risk Simulator) | 예측(Forecasting) | 최고의 우도 (Maximum Likelihood)**를 선택하여 주십시오.
- 하기의 리스크(Figure 3.20 를 참조)에서 종속 변수를 선택하고, **OK** 를 클릭하여 모델과 레포트를 실행하여 주십시오.

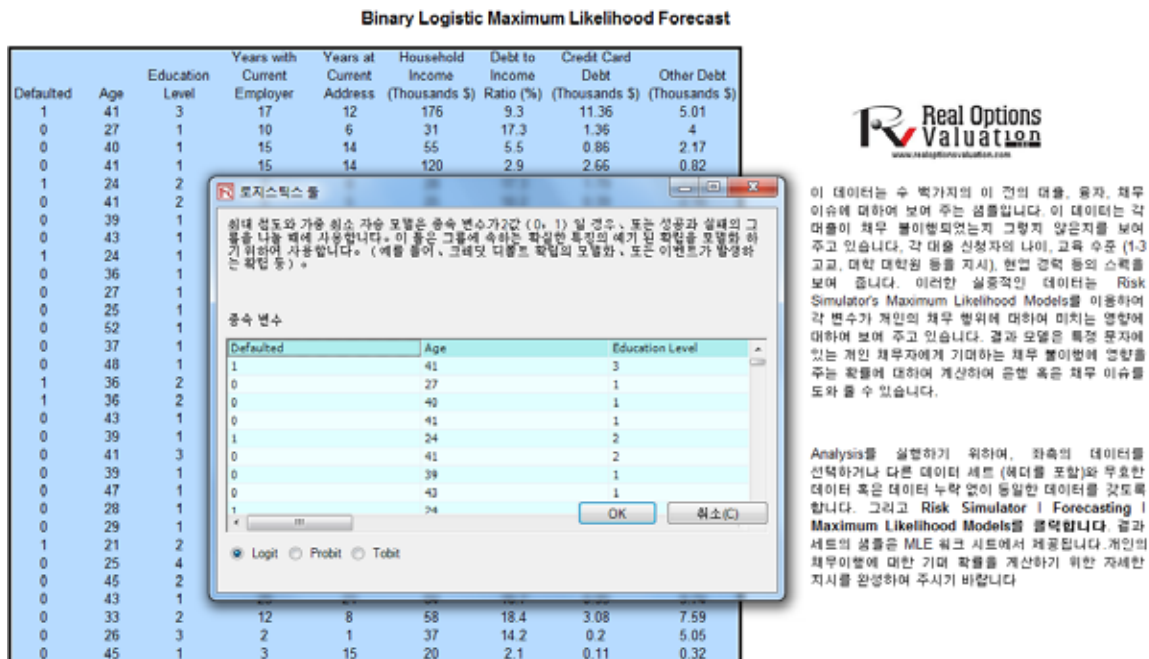


Figure 3.20 최고 우도의 모델

Spline (Cubic Spline Interpolation and Extrapolation/입방 스플라인)

이론:

때때로, 시계열 데이터의 설정으로 차손값이 있습니다. 예를 들어, 1 년에서 3 년의 이율은 존재하고, 5 년에서 8 년까지 이어져, 이후에는 10 년밖에 없는 경우를 상정하고 있습니다. 스플라인의 곡선은, 실재하는 데이터에 기초하여 잃어버린 년도의 이율을 투입하는 것에 사용됩니다. 따라서 스플라인 곡선은, 현재의 데이터의 시간 주기를 넘어서 미래의 시간 주기의 값을 예측, 또는 외삽하기 위하여 사용됩니다. 데이터는 선형, 및 비선형에 있는 것이 가능합니다. Figure 3.21 는, 어떻게 하여 입방 스플라인이 실행되는가를 표시합니다. 알려져 있는 X 값은, 차트의 x-축상의 값 (이 예증에서는, 알려져 있는 년도의 이율입니다) 를 표시하고, 알려져 있는 Y 값은, y-축상에서의 값 (이 케이스에서는 알려져 있는 이율) 을 표시하고 있습니다.

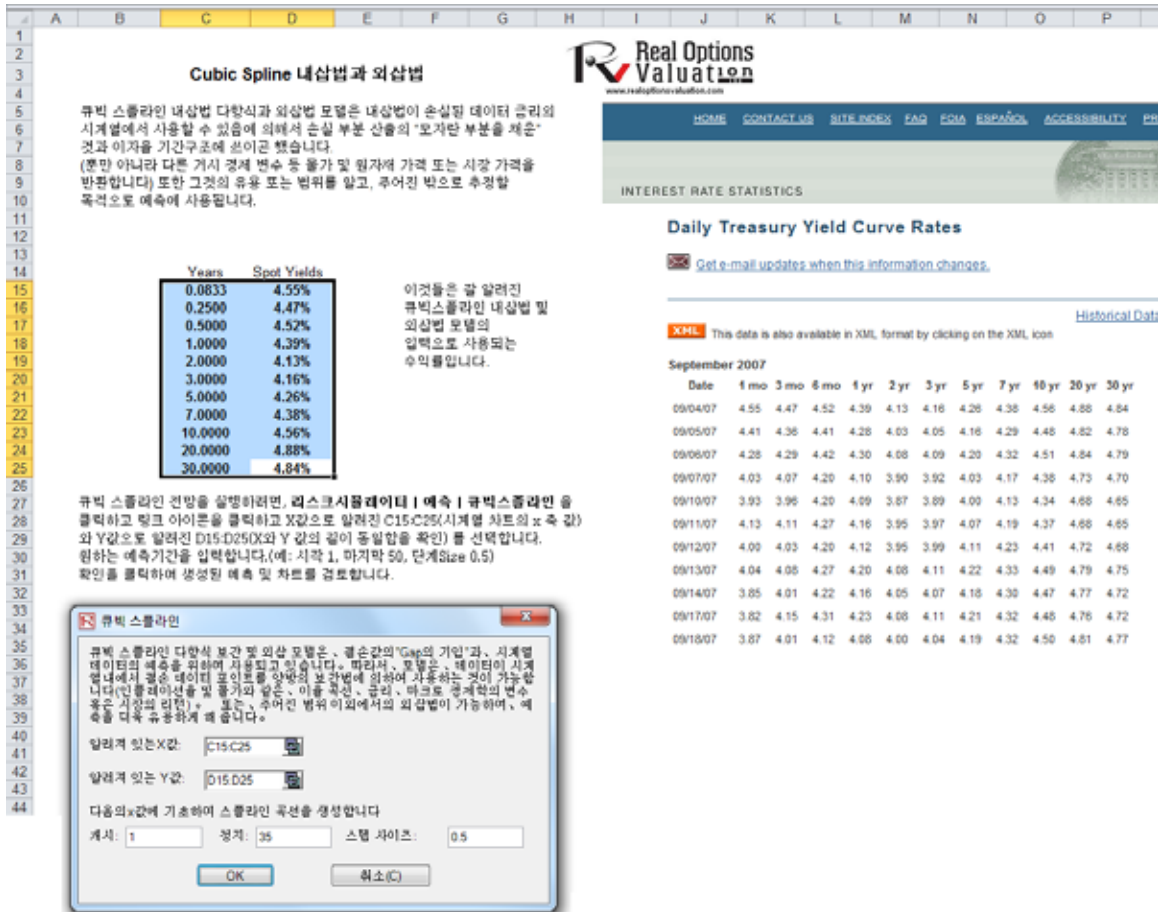


Figure 3.21 Cubic Spline Module

순서:

1. Excel 을 기동하고, 예증 파일의 **고도 예측 모델, 입방 스플라인** 워크 시트를 열고, 헤더를 포함한 데이터 설정을 선택하고, **리스크 시뮬레이터 (Risk Simulator) | 예측(Forecasting) | 입방 스플라인(Cubic Spline)**을 선택하여 주십시오.
2. 링크 아이콘을 클릭하고, 알려져 있는 X 값과 Y 값 (Figure 3.21 을 예증으로 참조하여 주십시오) 을 클릭하고, 투입, 및 외삽을 위하여 필요한 초기값, 정지값, 이것의 값 사이의 스텝 와이즈를 기입하고, **OK** 를 클릭하고 모델과 레포트를 실행하여 주십시오.

4.

이 장에서는, 리스크 시뮬레이터의 사용에 속하는 최적화 과정과 방법의 상세를 참조합니다. 여기에서는, 정적에 대하여 다이나믹, 그리고 확률적 최적화와 같이, 연속에 대하여 이산적 정수의 최적화 포함된 방법을 사용 합니다.

몇 개의 알고리즘은 최적화를 실행하기 위하여 존재하고, 최적화가 몬테카를로 시뮬레이션과 같이 실행되는 경우, 여러 가지 과정이 나타납니다. 리스크 시뮬레이터에서는, 세 개의 다른 최적화의 과정과, 다른 결정 변수의 타입과 같은 최적화 타입이 있습니다. 예를 들어, 리스크 시뮬레이터는, **정수 결정 변수** (예, 1, 2, 3, 4 혹은, 1.5, 2.5, 3.5 등), **2 직결정 변수** (네·아니오의 판단을 위한 1 과 0)과, **혼합결정변수** (양방향과 정수에서 연속적인 변수)와 같이, **연속 결정 변수** (1.2535, 0.2215 등)을 움직이는 것이 가능합니다. 그 이전에, 리스크 시뮬레이터는, **비선형적 최적화** (예, 목적과 규제 범위가 비선형적인 기능과 방식의 혼합과 같이, 선형의 혼합을 가졌을 때) 과 같이, **선형적 최적화** (예, 목적과 제약의 양방이 모두 선형적인 방정식이나 관수인 경우) 를 행할 수 있습니다.

최적화 프로세스에 관하여, 리스크 시뮬레이터는, **이산적 최적화**를 실행하는 것으로 사용할 수 있습니다. 최적화는, 이산적, 또는 정적 모델상에서, 시뮬레이션 없이 실행할 수 있습니다. 즉, 모델의 모든 입력은 정적으로 무변환이라는 것입니다. 이 최적화의 타입, 불확실성이 없이, 모델이 알려져 있으면 가정되는 때에 적용됩니다. 따라서, 최적의 포트 폴리오를 선택하여, 이것에 적합한 최적의 결정 변수의 변동을 내기 위하여, 고도로 최적화 과정을 행하기 전에, 또는 이산적인 최적화를 실행하는 것이 가능합니다. 예를 들어, 확률적으로 최적화 문제를 실행하기 전에, 실재하는 해결이 있는가를 보기 위하여, 연장된 분석을 적용하기 전에 이산적 최적화를 먼저 실행합니다.

다음에, 최적화가 몬테카를로 시뮬레이션과 같이 적용되는 때에, **다이나믹 한 최적화**를 사용합니다. **시뮬레이션·최적화**의 명칭으로도 알려져 있습니다.

이것은, 시뮬레이션이 먼저 실행되고, 그 결과는 Excel 모델에 적용되어, 최적화는 시뮬레이션 된 값에 적용됩니다. 즉, 시뮬레이션은 N 시행 실행되고, 그 후에, 최적화 과정은, 최적화 결과를 얻기 전에, 또는 실행 불가능한 설정이 나오기 전에 M 회 반복됩니다. 이것은, 리스크 시뮬레이터의 최적화 모듈의 사용에 의하여 어떠한 예측과 가정 통계를 사용하고, 시뮬레이션의 실행후에 바꿀 수 있으나, 선택이 가능합니다. 이 후에, 예측 통계가 최적화 과정상에서 적용됩니다. 이 어프로치는, 복수의 상호 작용이 되는 가정과 예측이 있는 큰 모델인 경우에, 그리고, 최적화에서 몇 개의 예측 통계가 필요할 때의 사용이 최적입니다. 예를 들어, 가정, 및 예측의 표준 오차가 최적화 모델로써 필요한

경우 (예, 평균을 포트폴리오의 표준 편차로 나눈 최적화 문제와 자산의 배당 샤프 지수) , 이 어프로치가 사용됩니다.

한편, 확률적 최적화의 과정은, 다이나믹 최적화의 과정과 비슷합니다만, 차이는 모든 다이나믹한 최적화의 과정은 **7회** 반복된다는 것입니다. 이것은, 시행 회수 N 을 가진 시뮬레이션이 실행되고, 그 후에, 최적의 결과를 얻기 위한 최적화가 반복수, M 실행됩니다. 이 후에, 과정은 7회 복제됩니다. 결과로서 7값과 같이 각 결정 변수의 예측 차트가 표시됩니다. 즉, 시뮬레이션이 실행되고, 예측, 또는 가정 통계는, 최적의 결정 변수의 할당을 내기 위하여 최적화 모델로 사용됩니다. 이 후에, 다른 시뮬레이션이 실행됩니다, 다른 예측 시계를 생성하고, 이것의 갱신 된 값은 다시 최적화 됩니다. 따라서, 최종 결정 변수는, 최적의 결정 변수의 범위를 표시하는 각자의 예측 차트를 가지고 있습니다. 예를 들어, 다이나믹한 최적화 과정에서 단일 · 포인트의 추정을 얻는 대신에, 결정 변수의 분포를 얻는 것이 가능, 따라서 각 결정 변수의 최적의 값의 범위를 얻는 것이 가능합니다. 즉 확률 과정으로서도 알려져 있습니다.

최후에, 유효 프론티어 최적화 과정이 최적화에 있어서의 Marginal 증분 및 웨도우 가격의 개념이 응용됩니다. 즉, 제약의 어느쪽인가 조금씩 느슨한 경우, 최적화의 결과에 무엇인가 일어나지 않겠습니까? 예를 들어 말하면, \$1,000,000 로서 예산의 제약이 설정되어 있는가 어떤가 입니다. 만약, 제약이 \$1,500,000, 또는 \$2,000,000 의 경우, 최적의 결정과 포트 폴리오의 결과에 무엇이 일어나지 않겠습니까? 이것은, 투자 파이낸스에서의 Markowitz 의 유효적 프론티어로, 포트 폴리오의 표준 편차가 느슨하게 증분하는 경우, 어떠한 부가적인 리턴을 포트 폴리오를 생성하는 것은 어떨습니까? 이 과정은, 한 개의 제약을 변환하는 것이 가능, 각 변환에 의하여 시뮬레이션과 최적화의 과정이 실행되는 것 이외에는, 다이나믹한 최적화의 과정에 유의하고 있습니다. 이 과정은, 리스크 시뮬레이션을 사용하고 수동적으로 적용하는 것이 최적입니다. 이것은 다이나믹, 및 확률적 최적화를 실행하는 것으로, 이 후에, 제약과 같이 다른 최적화를 재실행하고, 이 과정을 몇번이고 반복하는 것입니다. 이 수동적인 과정은, 부가적인 분석의 가치가 있으나, 또는 목적과 결정 변수의 유의한 변환을 얻기 위하여, 제약에서의 최저한의 증가가 어느 정도 떨어져 있지 않으면 안되는지, 제약의 변환에 의하여 결과가 유의하고 있으나, 상이한가를 정하는 분석과 같이 중요합니다.

한 개의 항목은 고찰의 가치가 있습니다. 확률적인 최적화를 실행하는 다른 소프트웨어가 있으나, 실용적이지 않습니다. 예를 들어, 시뮬레이션의 실행 후, 최적화 과정의 1 회의 반복이 생성되고, 다른 시뮬레이션이 실행된 후, 2 회째의 최적화의 반복이 실행되기 위하여, 시간과 자원의 낭비입니다.

이것은, 최적화에서는 , 최적의 결과를 얻기 위하여, 복수의 반복 (복수에서 천번의 반복을 불러옴) 이 필요한 엄격한 알고리즘의 설정을 통하여 모델에 도입되었습니다. 따라서, 한 개의 반복의 생성은, 시간과 자원의 낭비입니다. 같은

포트폴리오는, 몇 시간도 필요하기 전에 기술되는 어프로치를 사용하는 것에 의하여도, 리스크 시뮬레이터를 통하여 몇 분안에 해결하는 것이 가능합니다. 시뮬레이션 · 최적화의 어프로치는 일반적으로 나쁜 결과를 불러오지만, 확률적 최적화의 어프로치에서는 그렇지 않은 않습니다. 모델에 최적화를 적용할 때에는, 주의하여 주시기 바랍니다.

다음에 두 개의 최적화의 문제가 있습니다. 하나는, 연속적인 결정 변수를 사용하고, 다른 하나는, 이산적인 정수의 결정 변수를 사용하고 있습니다. 어느쪽인가의 모델에서, 이산적 최적화, 다이나믹한 최적화, 확률적인 최적화가 적용되고, 또 유효 프론티어와 쉐도의 가격을 적용할 수 있습니다. 이것의 어프로치의 어느쪽인가 두개의 예증에 사용 가능합니다. 따라서, 간결하게, 모델의 설정만이 표시되고, 사용자를 위하여 어떠한 최적화를 사용하는가의 선택이 가능합니다. 또는, 연속적 모델은, 정수의 최적화의 두 번째의 사례가 선형적 최적화 모델의 사례 (이것의 목적과 모든 제약은 선형적) 에 대하여, 비선형 최적화의 어프로치(이것은, 계산되는 포트 폴리오의 리스크는 비선형 공식을 가지고, 목적은, 포트 폴리오의 리턴의 비선형 공식을 포트폴리오의 리스크로 나눈 것입니다)를 사용합니다. 따라서, 이것의 두 개의 사례는, 전에 기술된 모든 과정이 포함됩니다.

Figure 4.1 는, 연속적인 최적화 모델의 샘플을 표시하고 있습니다. 여기에서 사용하고있는 사례는, **연속적 최적화**를 사용하고, 이 파일은 스타트 메뉴가 있습니다. **스타트 (Start) | 리얼 옵션즈 밸류에이션 (Real Options Valuation) | 리스크 시뮬레이터(Risk Simulator) | 예증 (Examples)** , 또는 직접 **리스크 시뮬레이터 Risk Simulator) | 예증 모델(Example Models)**을 통하여 접근하여 주십시오. 이 사례에서는, 10 개의 다른 자산 클래스 (예, 여러 가지 타입의 투자 신탁, 주식, 및 자산) 이 있어서, 아이디어로서는, 무엇보다도 유효한 할당을 가지고 있는 포트폴리오는 무엇보다도 최적의 이익을 얻는 것입니다. 이것은, 가능한 최적의 포트폴리오의 리턴을 생성하기 위하여는, 각 자산 클래스에 고정 리스크를 부여 할 수 있습니다. 최적화의 컨셉트를 해석을 하기 위하여, 어떤 식으로 하여 최적화의 과정이 보다 좋은 방법으로 적용될 것인가를 위하여 샘플 모델을 가지고 깊이 공부하지 않으면 안됩니다. 모델은, 10 개의 자산 클래스를 표시하고, 이것의 각 자산은 독자의 연간 리턴과 연간 예측 변동률을 가지고 있습니다. 이것의 리턴과 리스크의 측정은, 다른 자산 클래스와 일관적으로 비교될 수 있도록 연간 단위로 표시되어 있습니다. 리턴은, 비교 가능한 리턴의 기하학적 평균을 사용하고 계산되고 있는 것에 대하여, 리스크는, 대수적으로 비교 가능한 주식 리턴의 어프로치를 사용하고 계산됩니다. 주식 및 자산 클래스상에서의 연간 리턴과 연간 예측 변동률의 계산의 상세에서는, 이 장의 부록을 참조하여 주십시오.

	A	B	C	D	E	F	G	H	I	J	K	L
1												
2												
3												
4												
5												
6		자산 등급 종류	연간 수익	변동성 리스크	할당 가중치	최소 요구 할당	최대 요구 할당	리스크율 수익	수익순위 (Hi-Lo)	리스크 순위(Lo-Hi)	리스크순위 수익 (Hi-Lo)	할당순위 (Hi-Lo)
7		Asset Class 1	10.54%	12.36%	10.00%	5.00%	35.00%	0.8524	9	2	7	1
8		Asset Class 2	11.25%	16.23%	10.00%	5.00%	35.00%	0.6929	7	8	10	1
9		Asset Class 3	11.84%	15.64%	10.00%	5.00%	35.00%	0.7570	6	7	9	1
10		Asset Class 4	10.64%	12.35%	10.00%	5.00%	35.00%	0.8615	8	1	5	1
11		Asset Class 5	13.25%	13.28%	10.00%	5.00%	35.00%	0.9977	5	4	2	1
12		Asset Class 6	14.21%	14.39%	10.00%	5.00%	35.00%	0.9875	3	6	3	1
13		Asset Class 7	15.53%	14.25%	10.00%	5.00%	35.00%	1.0898	1	5	1	1
14		Asset Class 8	14.95%	16.44%	10.00%	5.00%	35.00%	0.9094	2	9	4	1
15		Asset Class 9	14.16%	16.50%	10.00%	5.00%	35.00%	0.8584	4	10	6	1
16		Asset Class 10	10.06%	12.50%	10.00%	5.00%	35.00%	0.8045	10	3	8	1
17		포트폴리오 관계	12.6419%	4.58%	100.00%							
18		리스크율의 수익	2.7596									
19												
20												
21												
22		목적:										
23		결정변수:										
24		결정변수에 대한 제한:										
25		제약:										
26		추가 사항:										
27												
28												
29												
30												
31												

Figure 4.1 연속 최적화의 모델

컬럼 E 에서의 할당 중량은, 종합적인 중량이 100% (셀 E17)로서 제약되는 것으로 가정, 그리고 돌릴 필요가 있는 변수, 결정 변수를 유지하고 있습니다.

일반적으로, 최적화를 시작하는 것에 있어서, 이것의 셀을 균일한 값에 설정하지 않으면 안되고, 이 케이스에서는, 셀 E6 에서 E15 까지의 각 셀은 10%로 설정됩니다. 또한, 각 결정 변수는, 정해진 범위에서 특정한 제한을 가지고 있습니다. 이 예증에서는, 하한, 및 상한의 정해진 할당에서는, 컬럼 F 와 G 와 같이 5%와 35%입니다. 이것은 각 자산 클래스가, 독자의 할당의 한계를 가지고 있는 것을 표시합니다. 다음에, 컬럼 H 는, 리스크에 대하여 리턴 비율을 표시하고, 단순히 리턴의 백분위 수를 리스크의 백분위 수로 나눈 결과로, 값이 높을수록, 높은 이익을 가져오는 것을 표시하고 있습니다. 남아 있는 모델은 개인적인 자산 클래스의 순위를 리턴, 리스크, 리스크에 대한 리턴의 비율과 배분에 나누어 표시하고 있습니다. 즉, 이것의 순위표는, 한 눈에 어떠한 자산 클래스가 낮은 리스크를 가지고 있을까 혹은 높은 리턴을 가지고 있는 것을 표시하고 있습니다.

셀 C17 에 있는 포트 폴리오의 종합적인 리턴은, $SUMPRODUCT(C6:C15, E6:E15)$ 로, 이것은, 배분 중량과 각 자산 클래스의 연간 리턴의 적을 더한 것입니다. 즉, $R_p = \omega_A R_A + \omega_B R_B + \omega_C R_C + \omega_D R_D$ 로, R_p 는, 포트 폴리오의 리턴을 표시하고, $R_{A,B,C,D}$ 는, 프로젝트의 개별 리턴을 표시하고, $\omega_{A,B,C,D}$ 는, 각 중량이 각 프로젝트의 자산 할당입니다.

또는, 셀 D17 에 다각화 되어 있는 리스크의 포트

폴리오는,
$$\sigma_P = \sqrt{\sum_{i=1}^i \omega_i^2 \sigma_i^2 + \sum_{i=1}^n \sum_{j=1}^m 2\omega_i \omega_j \rho_{i,j} \sigma_i \sigma_j}$$
 에 의하여 계산되고

있습니다. 여기서는 $\rho_{i,j}$ 는, 자산 클래스 사이의 상관 관계에 해당하고, 따라서, 상호상관이 마이너스인 경우, 리스크의 다양화의 효과가 보여지고, 포트 폴리오의 리스크는 감소합니다. 단, 계산을 간단히 하기 위하여, 이 포트 폴리오의 리스크 계산을 통하여, 자산 클래스중의 상관을 제로로 고려하고 있으나, 먼저 기술 한 것과 같이, 리턴의 시뮬레이션의 적용 시에, 상관을 고려하여 주십시오. 따라서, 이것의 다른 자산 리턴 중의 정적 상관을 적용하는 대신에, 이것 자신의 시뮬레이션 가정의 상관을 적용하고, 시뮬레이션된 리턴의 값 사이에 더욱ダイナ믹한 관계가 나타납니다.

최후에, 리스크의 리턴에의 비율, 및 샤프율은, 포트 폴리오를 위하여 계산됩니다. 이 값은, 셀의 C18 로 표시되고, 이 최적화의 실행으로 최대화된 목적을 표시하고 있습니다. 마무리로, 이 예증 모델의 상세가 하기에 기술되어 있습니다.

목적: 리스크에의 최대의 리턴 비율(C18)
 결정 변수: 할당 중량 (E6:E15)
 결정 변수의 제한: 필요한 최소값과 최대값 (F6:G15)
 제약: 종합적인 할당의 덧셈은 100% (E17)

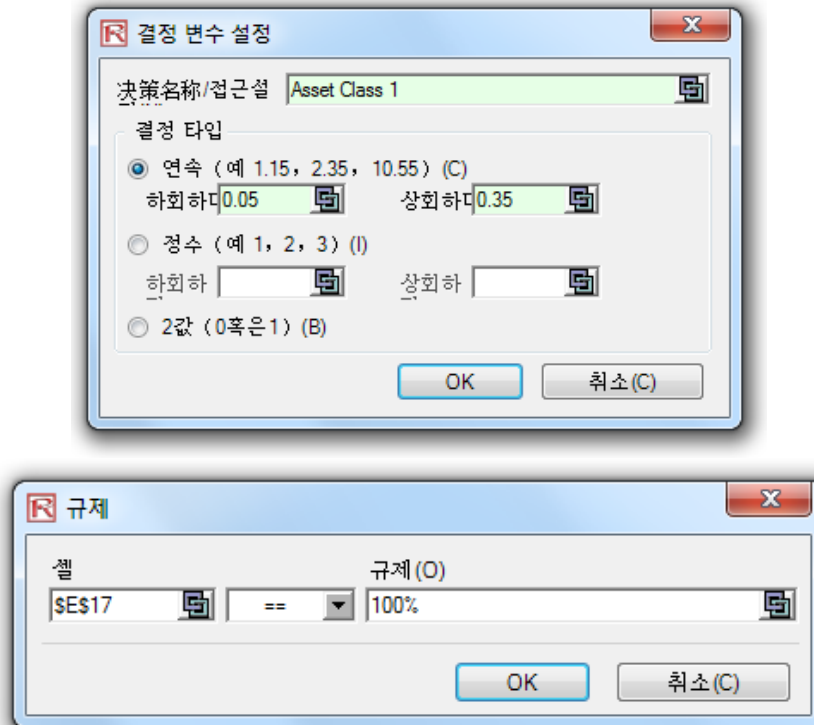
순서:

- 예증 파일을 열고, 신규 파일을 **리스크 시뮬레이터 (Risk Simulator) |신규 프로필 (New Profile)** 클릭하고 기동하여, 명칭을 붙여 주십시오.
- 최적화의 최초의 스텝으로서 결정 변수의 설정을 행하여 주십시오. 셀의 E6 를 선택하고, 최초의 결정 변수를 설정 (**리스크 시뮬레이터 (Risk Simulator) |최적화 (Optimization) |결정의 설정 (Set Decision)**)을 하여, 링크 아이콘을 클릭하고, 셀 F6 과 G6 의 하한과 상한의 값과 같이 명칭 셀(B6)을 선택하여 주십시오. 그 후에, 리스크 시뮬레이터의 카피를 사용하여, 셀 E6 의 결정 변수를 카피하여 남아 있는 셀의 E7 에서 E15 에 붙여 주십시오.
- 최적화의 두 번째의 스텝으로서, 제약을 설정하는 것입니다. 여기서는 하나의 제약밖에 없습니다. 이것은, 포트 폴리오의 종합적 할당의 덧셈은 100%가 되지 않으면 안되는 것을 의미합니다. 이를 위하여, **리스크 시뮬레이터 (Risk Simulator) | 최적화 (Optimization) | 제약 (Constraints) ... 그리고, 추가**를 선택하고 새로운 제약을 추가하여 주십시오. 그 후에, 셀 E17 을 선택하고, 100%에 동일(=)시켜 주십시오. 종료 후, OK 를 클릭하여 주십시오.
- 최적화의 최후의 스텝으로서, 목적 함수의 설정과 목적의 셀 C18 과 **리스크 시뮬레이터 (Risk Simulator) | 최적화 (Optimization) | 최적화의 실행 (Run Optimization)** 의 선택을 통하여 최적화를 실행시키고, 최적화 (정적 최적화, 다이내믹 최적화, 및 확률적 최적화) 의 선택을 하여

주십시오. 기동하는 것에 있어서, **정적 최적화**를 선택하여 주십시오. 목적 셀은 C18에 설정되어 있는가 어떤가를 확인한 후, 최대화를 선택하여 주십시오. 필요하다면, 제약과 결정 변수를 검토하는 것이 가능합니다. 또는, 정적 최적화를 실행하는데 있어서 OK를 클릭하여 주십시오.

최적화가 한번 완료한 후, 목적과 같이, 결정 변수를 원래 값으로 복귀시키기 위하여, **복귀**를 선택하거나, 최적화 된 결정 변수를 적용하기 위하여, **바꾸기**를 선택하여 주십시오. 일반적으로, 바꾸기는, 최적화의 실행후에 선택됩니다.

Figure 4.2는, 이것의 상기의 과정의 스텝을 표시하고 있습니다. 모델상의 리턴과 리스크(컬럼 C와 D)에 시뮬레이션 가정을 추가하는 것이 가능, 다이나믹한 최적화와 확률적 최적화를 부가 연습으로서 적용하는 것이 가능합니다.



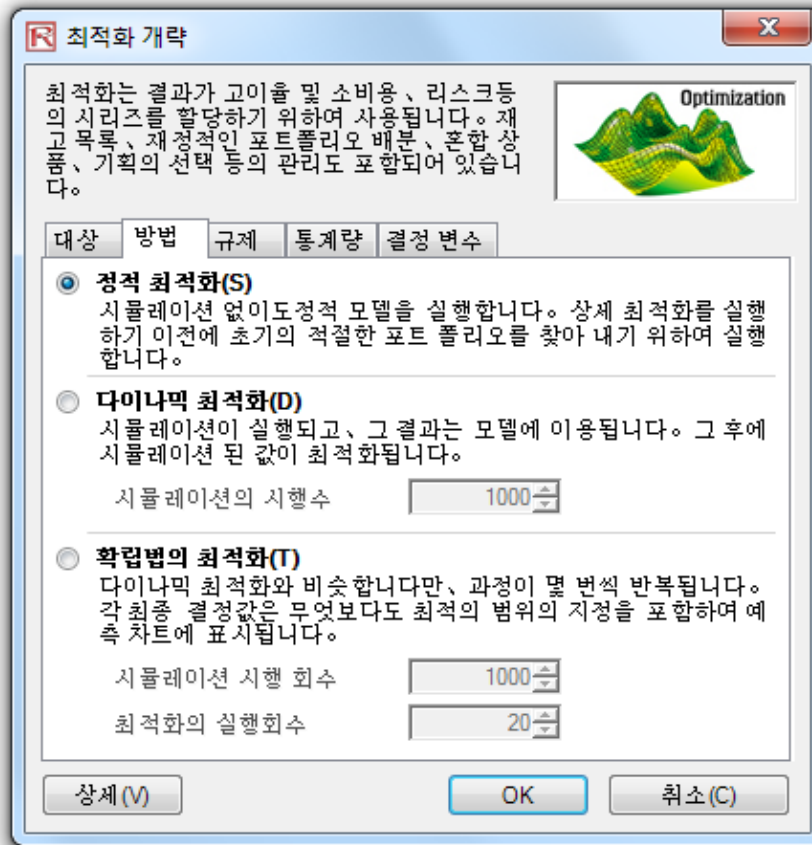


Figure 4.2 리스크 시뮬레이터에서의 연속 최적화 실행

결과의 해석:

최적화의 최종 결과는 Figure 4.3 에 표시되어 있어, 포트 폴리오를 위한 자산의 최적의 할당은 셀 E6:E15 에 있습니다. 이것은, 각 자산의 변동이 5%와 35%의 사이에서, 할당의 덧셈이 100%가 아니면 않된다는 것의 제한을 부여하고, 리스크에 리턴 비율의 최대화는 Figure 4.3 에 표시되어 있습니다.

이제까지의 결과와 최적화의 가정의 적용의 검토에서 다소 중요한 것은 메모하지 않으면 안됩니다.

- 바른 최적화의 실행은 이익, 및 리스크에 대한 리턴 샤프율을 최대화하는 것입니다.
- 만약, 그 대신에 종합 포트 폴리오의 리턴을 최대화 한 경우, 최적의 할당의 결과는 부족하고, 취득에 최적화는 필요 없습니다. 이것은, 낮은 8 개의 자산에 대하여 할당 5% (인정된 최소값) , 무엇보다도 리턴이 높은 자산에 대하여 35%(인정된 최대값) , 나머지(25%)는, 2 번째로 좋은 리턴을 가진 자산에 해당됩니다. 최적화는 필요하지 않습니다. 단, 이러한 방법으로 포트 폴리오를 배당하는 시기는, 포트 폴리오의 리턴 자신도 높게 되어 있으나, 리스크에의 리턴 비율을 최대화하는 시기에 비하여 리스크가 크게 높아집니다.
- 한편, 종합적인 포트 폴리오 리스크를 감소하는 것이 가능합니다만, 리턴도 낮아집니다.

Table 4.1 는, 최적화된 세 개의 다른 목적에서 된 결과를 표시하고 있습니다.

목적:	포트 폴리오의 리턴	포트 폴리오의 리스크	포트 폴리오의 리스크의 리턴 비율
리스크에의 리턴 비율의 최대화	12.69%	4.52%	2.8091
리턴의 최대화	13.97%	6.77%	2.0636
리스크의 최소화	12.38%	4.46%	2.7754

Table 4.1 최적화의 결과

표에서, 최적의 어프로치는 리스크에의 리턴 비유를 최대화하는 것에 있고, 이것은, 같은 양의 리스크에는, 이 할당은 무엇보다도 높은 리턴을 불러들입니다. 한편, 같은 양의 리턴은, 이 할당은 무엇보다도 낮은 리스크 확률을 부여 하여 줍니다. 이 이익, 및 리스크에의 리턴 비율의 어프로치는, 모던한 포트 폴리오 이론인 Markowitz 의 유효 프론티어의 초석입니다. 이것은, 종합적인 포트폴리오의 리스크 레벨을 제약하고, 시간이 흐름에 따라 이것을 증가하는 것으로, 다른 리스크 성질을 위한 여러 가지 유효한 포트 폴리오의 할당을 얻는 것이 가능합니다. 따라서, 다른 유효 포트폴리오의 배분은, 다른 리스크 선호를 가진 다른 개개인을 위하여 얻는 것이 가능합니다.

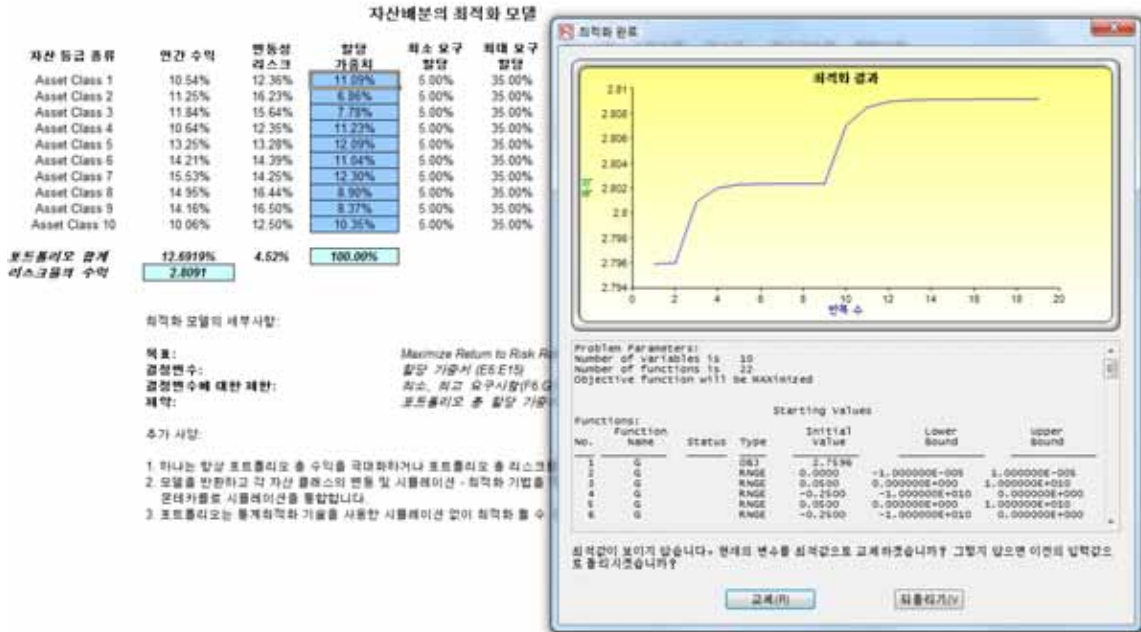


Figure 4.3 연속 최적화의 결과

때때로, 결정 변수는 연속적이 아닙니다만, 이산적 정수 (예, 0 과 1)입니다. 이것은, 이러한 최적화를 on-off 의 스위치, 또는 go/no-go 의 판단에 적용하는 것이 가능합니다. Figure 4.4 는, 20 의 프로젝트가 표시된 중에서의 프로젝트 선택 모델을 표시하고 있습니다. 여기에서의 예증은, **이산적인 최적화**파일을 사용하고, 스타트 메뉴에서 **스타트 (Start) | 리얼 옵션즈 밸류에이션(Real Options Valuation) | 리스크 시뮬레이터(Risk Simulator) | 예증(Examples)**을 선택하거나, **리스크 시뮬레이터 (Risk Simulator) | 예증 모델(Example Models)**을 클릭하고 직접 열어 주십시오. 전과 같이 각 프로젝트는, 독자의 리턴 (확장된 순수 현재 가치와 현재 가격을 위하여 ENPV DHD 와 NPV 를 가지고 있어, ENPV 는, 단순히 NPV 에 있는 전략적인 리얼 옵션의 값이 더해지는 것으로)로써, 실시의 비용, 리스크 등을 가지고 있습니다. 필요한 경우, 이 모델은 필요한 Full·Time 타임의 평등 가격(FTE) 、다른 여러 가지 기능의 자원, 그래서 이것의 부가 자원에 설정할 수 있는 제약을 포함하도록 수정하고 있습니다. 이 모델 중의 입력은 일반적으로, 다른 스프레트 시트 모델에 의하여 링크됩니다. 예를 들어, 각 프로젝트는, 독자의 DCF, 및 자산 모델의 리턴을 가지고 있습니다. 여기에서의 어플리케이션은, 어떤 예산 배분의 제약에 있어서 포트폴리오의 민감도를 최대화하는 것입니다. 이 모델의 여러 가지 버전을 작성할 수 있습니다. 예를 들어, 포트폴리오의 리턴을 최대화하거나, 리스크를 최소화하거나, 선택한 프로젝트의 종합적인 수는 6 을 넘지 않는 등의 제약을 추가합니다. 이것의 모든 아이템은, 존재하는 모델을 사용하고 실행하는 것이 가능합니다.

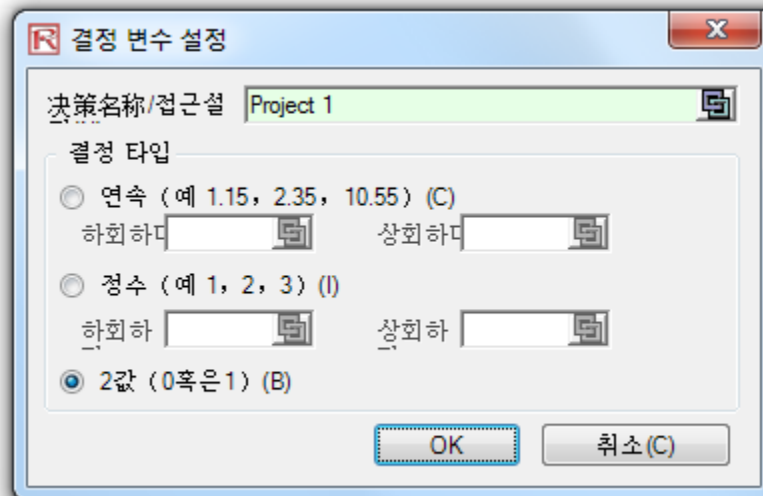
순서:

- 예증 파일을 열고, **리스크 시뮬레이터 (Risk Simulator) / 신규 프로파일 (New Profile)** 을 클릭하는 것으로 신규 프로파일을 기동하고 명칭을 부여하여 주십시오.
- 최적화의 최고의 스텝으로서 결정 변수의 설정을 행하여 주십시오. 셀의 J4 를 선택하고, 최초의 결정 변수를 설정 (**리스크 시뮬레이터 (Risk Simulator) / 최적화 (Optimization) / 결정의 설정 (Set Decision)**)을 하고, 링크 아이콘을 클릭하고, 명칭 셀(B4)을 붙여서, **두값**의 변수를 선택하여 주십시오. 그 후에, 리스크 시뮬레이터의 카피를 사용하고, 셀 J4 의 결정 변수를 카피하고, 남은 셀의 J5 에서 J15 에 붙여 주시기 바랍니다. 이것은, 여러가지의 결정 변수밖에 가지고 있지 않았을 때와, 각 결정 변수에 독특한 독자의 명칭을 붙이는 것이 가능한 경우에 최적의 방법입니다.
- 최적화의 두 번째의 스텝으로, 제약을 설정하는 것입니다. 여기서는 두 개의 제약이 있습니다. 이것은, 포트 폴리오의 종합적인 배당 예산은, \$5,000 이상이 아니면 안되고, 프로젝트의 종합수는 6 을 넘지 않으면 안되는 것을 의미하고 있습니다. 이를 위하여, **리스크 시뮬레이터 (Risk Simulator) / 최적화 (Optimization) / 제약 (Constraints) ... 그리고, 추가** 를 선택하고, 새로운 제약을 추가 하십시오. 그 후, 셀 D17 을 선택하고, 5,000 에 같이, 또는 보다 작게($\leq$)하여 주십시오. 셀 J17 ≤ 6 로 설정한 후, 반복하여 주십시오.
- 최적화의 최후의 스텝으로서, 목적 관수의 설정과 목적의 셀 C19 (또는 C17)과 **리스크 시뮬레이터 (Risk Simulator) / 최적화 (Optimization) | 목적의 설정 (Set Objective)** 의 선택을 통하여 최적화를 실행시키고, **리스크 시뮬레이터 (Risk Simulator) / 최적화 (Optimization) / 최적화의 실행 (Run Optimization)** 을 선택하고, 최적화 (정적 최적화, 다이나믹 최적화, 및 확률적 최적화) 의 선택을 하여 주십시오. 기동하는 것에 있어서, **정적 최적화**를 선택하여 주십시오. 목적 셀은, 샤프비율이 포트폴리오의 리스크에의 리턴비율이 어떠한가를 확인한 후, 최대화를 선택하여 주십시오. 필요하다면, 제약과 결정 변수를 검토하는 것이 가능합니다. 또는, 정적 최적화를 실행하는 것에 있어서 OK 를 클릭하여 주십시오.

Figure 4.5 는, 전에 기술된 과정을 표시하고 있습니다. 모델의 ENPV 와 리스크 (칼럼 C 와 F)상에 시뮬레이션 가정을 추가하는 것이 가능, 부가의 연습을 위하여, 다이나믹한 최적화와 확률적 최적화를 적용할 수 있습니다.

	A	B	C	D	E	F	G	H	I	J
1										
2										
3			ENPV	Cost	Risk \$	Risk %	Return to Risk Ratio	Profitability Index		Selection
4		Project 1	\$458.00	\$1,732.44	\$54.96	12.00%	8.33	1.26		1.0000
5		Project 2	\$1,954.00	\$859.00	\$1,914.92	98.00%	1.02	3.27		1.0000
6		Project 3	\$1,599.00	\$1,845.00	\$1,551.03	97.00%	1.03	1.87		1.0000
7		Project 4	\$2,251.00	\$1,645.00	\$1,012.95	45.00%	2.22	2.37		1.0000
8		Project 5	\$849.00	\$458.00	\$925.41	109.00%	0.92	2.85		1.0000
9		Project 6	\$758.00	\$52.00	\$560.92	74.00%	1.35	15.58		1.0000
10		Project 7	\$2,845.00	\$758.00	\$5,633.10	198.00%	0.51	4.75		1.0000
11		Project 8	\$1,235.00	\$115.00	\$926.25	75.00%	1.33	11.74		1.0000
12		Project 9	\$1,945.00	\$125.00	\$2,100.60	108.00%	0.93	16.56		1.0000
13		Project 10	\$2,250.00	\$458.00	\$1,912.50	85.00%	1.18	5.91		1.0000
14		Project 11	\$549.00	\$45.00	\$263.52	48.00%	2.08	13.20		1.0000
15		Project 12	\$525.00	\$105.00	\$309.75	59.00%	1.69	6.00		1.0000
16										
17		Total	\$17,218.00	\$8,197.44	\$7,007	40.70%				12.00
18		Goal:	MAX	<=\$5000						<=6
19		Sharpe Ratio	2.4573							
20										
21		ENPV is the expected NPV of each credit line or project, while Cost can be the total cost of								
22		administration as well as required capital holdings to cover the credit line, and Risk is the								
23		Coefficient of Variation of the credit line's ENPV.								

Figure 4.4 리스크 시뮬레이터에서의 이산계 정수의 최적화



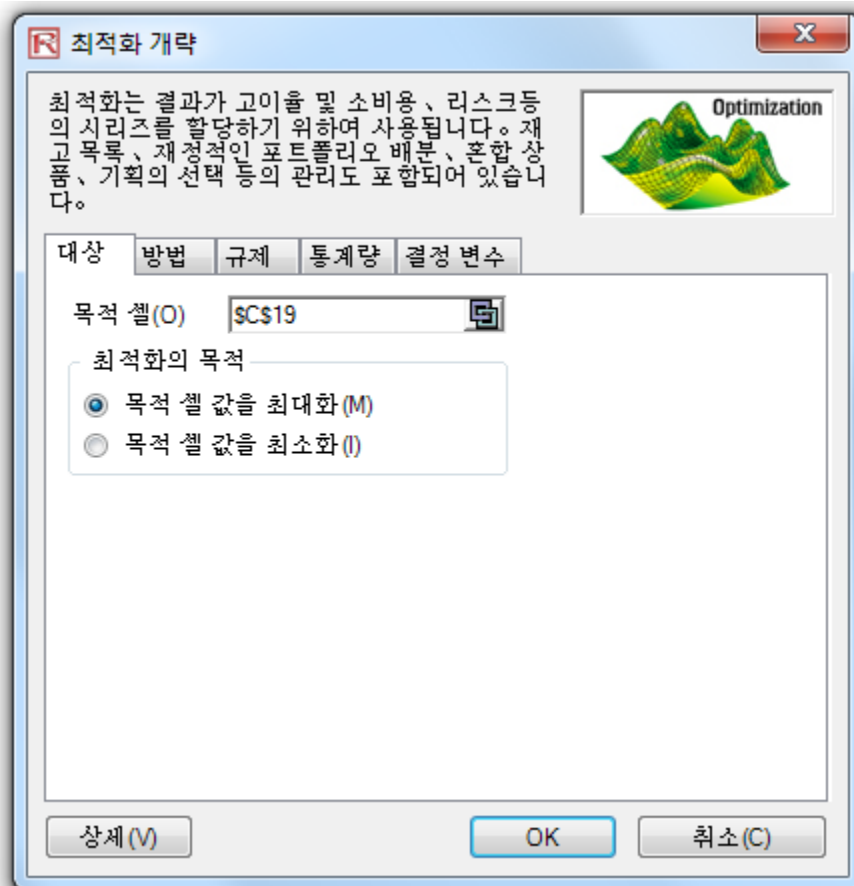
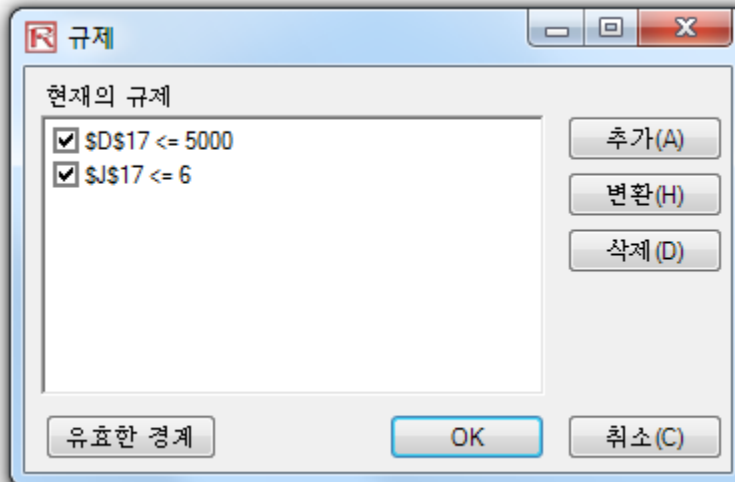


Figure 4.5 리스크 시뮬레이터에서의 이산계 정수의 최적화 실행

결과의 해석:

Figure 4.6 는, 샤프비율 최대화 하는 프로젝트의 최적의 선택 샘플을 표시하고 있습니다. 한편, 한 개는 항상 종합 수입을 최대화 할 수 있으나, 이것은 전술한 것과 같이 빈약한 과정으로, 무엇보다 높은 리턴의 프로젝트를 선택하는 것을 포함하고, 제약 예산을 초과할 때까지, 또는 돈을 사용하기 까지 리스크를 따라 갑니다. 이 실행은, 높은 이익을 가지고 오는 프로젝트는 일반적으로 높은 리스크를 가지고 있는 것으로, 이론적으로 바람직하지 않은 프로젝트를 불러 옵니다. 원하는 경우, 리스크의 값과 ENPV 에 가정을 추가하는 것으로, 확률, 또는 다이나믹한 최적화를 사용하는 것으로 최적화를 복제하는 것이 가능합니다.

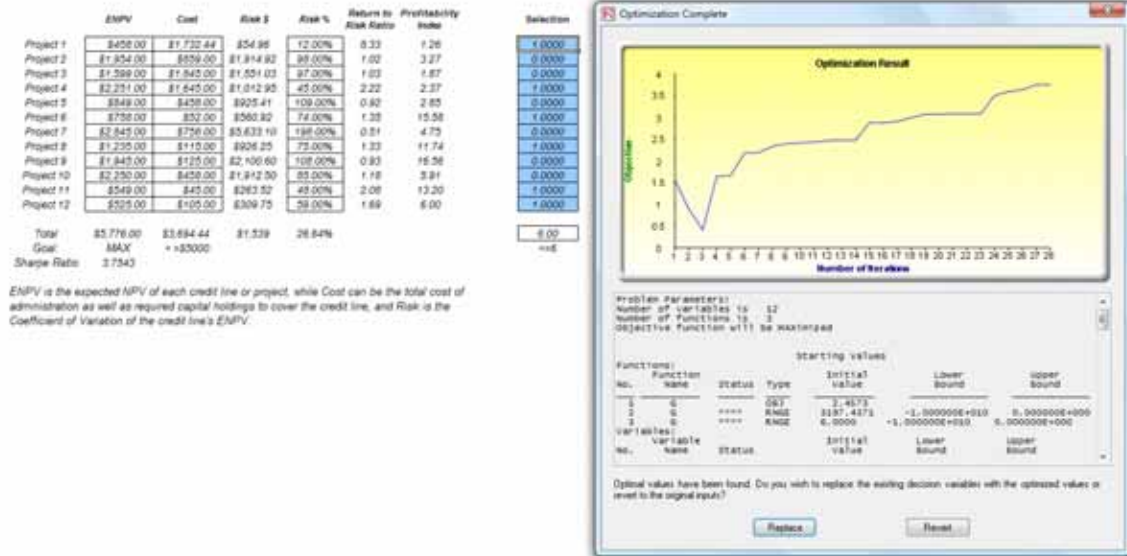


Figure 4.6 샤프 비율을 최대화 하는 최적화 프로젝트의 선택

활동하고 있는 최적화의 부가 예증에서는, 적분된 리스크 분석의 제 11 장의 케이스 스터디의 리얼 옵션 분석: *통과 기법, 제 2 판* (Wiley Finance, 2005)을 참조하여 주십시오. 이 케이스는, 어떻게 하여 유효한 프론티어가 생성될 수 있을까, 어떻게 하여 예측, 시뮬레이션, 최적화와 리얼 옵션이 유연하게 분석 과정중에 혼합할 수 있는가를 표시하고 있습니다.

5.

이 장에서는, 리스크 시뮬레이터의 분석 툴과 함께 취급하겠습니다. 이 분석 툴은 리스크 시뮬레이터의 소프트웨어의 예증의 어플리케이션을 통하여 볼 수 있습니다. 스텝 별로 표시되어 있는 일러스트와 함께 보시기 바랍니다. 리스크 분석 분야에서의 분석가의 일로서 이 툴은 매우 중요합니다. 적용 가능한 각 툴은 이 장에서 자세히 기술 되어 있습니다.

이론:

하나의 강력한 시뮬레이션 툴인 토네이도 분석은, 모델의 결과에의 각 변수의 정적인 영향을 얻을 수 있습니다. 이것은 모델내의 각 변수를 초기값에 자동적으로 부여하고 모델의 예측상의 변수를 얻어서, 결과값을 가장 유의한 것에서 그렇지 않은 순서대로 랭크합니다. Figures 5.1 에서 5.6 는, 토네이도 분석의 어플리케이션을 표시하고 있습니다. 예를 들어, Figure 5.1 은, 심플하게 나눈 캐쉬 플로의 모델로서, 모델의 입력 가정이 표시되고 있습니다. 여기에서의 의문은, 모델의 결과에 무엇보다 영향을 주는 크리티컬한 성공의 드라이버는 무엇인가? 라는 것입니다. 이것은, 어느 정도 정말로 \$96.63 의 순수 현재 가치(NPV)를 움직이는가, 또는, 어느 정도 입력 변수가 첫번째값에 영향을 주는가 ?라는 것입니다.

토네이도 툴은, **리스크 시뮬레이터 (Risk Simulator) 툴 (Tools) 토네이도 분석 (Tornado Analysis)** 을 통하여 얻는 것이 가능합니다. 또한, 최초의 예증에 이어서, 예증 필터에서 **토네이도와 감도 차트 (선택)** 필터를 엽니다. Figure 5.2 는 NPV 가 포함된 셀 G6 가 분석되 목적의 결과로서 선택된 샘플 모델이 표시되어 있습니다. 모델의 목적 셀의 선택은, 토네이도 차트를 작성하기 위하여 사용됩니다. 선택으로는, 모델의 결과에 영향을 주는 모든 입력과 종간의 변수입니다. 예를 들어, 모델은 $A = B + C$ 로, 및 $C = D + E$ 로, B, D, 와 E 는 A 의 선택(C 는, 선택이 아니라 종간의 계산된 값의 경우입니다)입니다. Figure 5.2 는, 목적의 결과의 추정을 위하여, 각 선택의 변수의 테스트 범위를 표시하고 있습니다. 만약에, 선택의 변수가 심플한 입력인 경우, 선택된 범위 (예, 디폴트는 $\pm 10\%$) 에 기초한 검정의 범위는, 심플한 합니다. 각 선택의 변수는, 필요한 경우, 여러가지 백분위 수로 교란될 수 있습니다. 예기 된 값의 주위의 작은 혼란보다 보다 극단의 값이 검정하기 쉬운 것처럼 폭넓은 범위는 중요합니다. 특정의 상황에서는, 극단값은 크고, 작고, 또는 언밸런스는 임팩트 (예, 비선형은, 변수의 대소값에 대한 규모의 클립과 경제 스케일의 증가, 또는 감소가 있을 경우에 발생합니다) 를 가지고 있을 것임으로, 폭넓은 범위만 이 비선형의 임팩트를 얻는 것이 가능합니다.

	A	B	C	D	E	F	G
1							
2							
3							
4		기본 년도	2005		PV 순이익 합계		\$1,896.63
5		조정된 시장리스크 할인율	15.00%		PV 투자 합계		\$1,800.00
6		개인 리스크 할인율	5.00%		순현재가치		\$96.63
7		연간 매출증가율	2.00%		회귀 내부율		18.80%
8		금액약화를	5.00%		투자액의 회귀		5.37%
9		실제 세율	40.00%				
10							
11			2005	2006	2007	2008	2009
12		제품A 평균값	\$10.00	\$9.50	\$9.03	\$8.57	\$8.15
13		제품B 평균값	\$12.25	\$11.64	\$11.06	\$10.50	\$9.98
14		제품C 평균값	\$15.15	\$14.39	\$13.67	\$12.99	\$12.34
15		제품A 수량	50.00	51.00	52.02	53.06	54.12
16		제품B 수량	35.00	35.70	36.41	37.14	37.89
17		제품C 수량	20.00	20.40	20.81	21.22	21.65
18		총 수익	\$1,231.75	\$1,193.57	\$1,156.57	\$1,120.71	\$1,085.97
19		상품판매가격	\$184.76	\$179.03	\$173.48	\$168.11	\$162.90
20		매출 총 이익	\$1,046.99	\$1,014.53	\$983.08	\$952.60	\$923.07
21		운영비용	\$157.50	\$160.65	\$163.86	\$167.14	\$170.48
22		판매비와 비용	\$15.75	\$16.07	\$16.39	\$16.71	\$17.05
23		영업이익 (EBITDA)	\$873.74	\$837.82	\$802.83	\$768.75	\$735.54
24		감가상각비	\$10.00	\$10.00	\$10.00	\$10.00	\$10.00
25		상각	\$3.00	\$3.00	\$3.00	\$3.00	\$3.00
26		EBIT	\$860.74	\$824.82	\$789.83	\$755.75	\$722.54
27		이자지급	\$2.00	\$2.00	\$2.00	\$2.00	\$2.00
28		EBT	\$858.74	\$822.82	\$787.83	\$753.75	\$720.54
29		세금	\$343.50	\$329.13	\$315.13	\$301.50	\$288.22
30		당기 순이익	\$515.24	\$493.69	\$472.70	\$452.25	\$432.33
31		감가상각비	\$13.00	\$13.00	\$13.00	\$13.00	\$13.00
32		순자산 변동	\$0.00	\$0.00	\$0.00	\$0.00	\$0.00
33		자본지출	\$0.00	\$0.00	\$0.00	\$0.00	\$0.00
34		Free Cash Flow	\$528.24	\$506.69	\$485.70	\$465.25	\$445.33
35							
36		투자	\$1,800.00				
37							

Figure 5.1: 모범 모델

순서:

- Excel 모델(예, 이 예증에서는 셀의 G6 가 선택되고 있습니다)에서 단일의 결과 셀(예, 공식, 또는 방식을 포함한 셀)을 선택하여 주십시오.
- 리스크 시뮬레이터 (Risk Simulator) / 툴(Tools) / 토네이도 분석(Tornado Analysis)**을 선택하여 주십시오.
- 선례를 확인하고, 적절한 명칭을 붙여서 고치도록 합시다(선례를 간단히 명칭으로 붙이고 토네이도 와 스파이더 차트가 보기 쉽게 됩니다. 그후에, OK 를 클릭하여 주십시오.



Figure 5.2 – 토네이도 분석의 실행

결과의 해석:

Figure 5.3 은, 투자 자산은 NPV 에 영향을 주는 것을 표시하고, 세율, 평균 판매 가격, 요구된 상품 라인의 양등에 의하여 따라가고, 토네이도 분석의 결과의 레포트를 표시합니다. 레포트는 4 개의 다른 요소를 포함하고 있습니다.

- ① 실행된 과정의 리스크의 통계적 개략
- ② 감도표(Figure 5.4)는, 초기의 NPV 의 기초값은 어떻게 각 입력 (예, 투자는 상회+10% 의 \$1,800 에서 \$1,980 로 변환되고, 하회는-10% 의 \$1,800 에서 \$1,620 로 변환)이 변환되는 가를 표시하고 있습니다. 결과가 된 NPV 의 상회, 그리고 하회는, -\$83.37 와 \$276.63 로, \$360 의 종합 변수로, NPV 상의 무엇보다도 높은 영향을 주는 변수가 됩니다. 선례 변수는 높은 영향을 주는 것에서 그렇지 않은 순서대로 랭크합니다.
- ③ 스파이더 차트는(Figure 5.5), 이것의 효과의 그래프를 표시하고 있습니다. Y-축은 NPV 의 목적치인 것에 대하여, x-축은, 각 선례의 값(중심 포인트는, 0%의 변환으로서 96.63 의 기본적인 케이스값에서 각 선례의

기본적인 값입니다)사의 백분위수의 변환을 표시하고 있습니다. 플러스 경사의 선은, 긍정의 관계, 및 효과를 표시하고, 마이너스 경향의 선은, 마이너스 관계를 표시합니다(예, 투자는 마이너스로서 슬로프되고, 투자의 레벨이 높을수록, NPV 는 낮게 됩니다). 슬로프의 절대값은 효과의 크기(스텝 라인은 NPV 상에서의 선례의 x-축의 변환에 부여된 y-축에서 가장 높은 영향을 표시하고 있습니다)를 표시하고 있습니다).

- ⑩ 토네이도는, 이 그래프를 다른 방법으로 표시하고, 무엇보다 높은 영향을 주는 것이 최초로 표시되고 있습니다. x-축은, NPV 의 값과 기본적이 케이스의 조건의 차트의 중심입니다. 차트의 녹색은, 플러스의 효과를 표시하고 있는 것에 대하여, 붉은 부분은, 마이너스의 효과를 표시하고 있습니다. 따라서, 투자를 위하여 우측의 붉은 색은 높은 NPV 상에서의 마이너스 투자를 표시하고 있습니다. 즉, 투자 자본의 NPV 는, 마이너스 관계에서 상관되는 것을 의미합니다. 역으로, 가격과 상품 A 에서 C (차트의 좌측에 이것의 녹색 부분이 있습니다) 의 양에서는 대립적인 결과가 보여집니다.

있습니다. 결과는 작은 영향을 주고, 불확실하다고 정해진 변수의 시뮬레이션에 시간을 뺏기지 마십시오. 토네이도 차트는, 이것의 크리티컬한 성공 드라이버를 빨리 그리고 간단히 보여주도록 도와 줍니다. 이 예증을 따라가면, 만약, 요구된 투자와 유효 세율의 양방이 사전에 알려져 있어서, 무변화라고 가정 하면, 가격과 양은 시뮬레이션되지 않으면 됩니다.

선례의 셀	기준 가치: 96.6261638553219			입력 변환		
	하측의 출력	상측의 출력	유리한 범위	하측의 입력	상측의 입력	베이스 케이스의 값
C36: 투자	276.6261639	-83.3738361	360.00	\$1,620.00	\$1,980.00	\$1,800.00
C9: 실제 세율	219.7269269	-26.4745992	246.20	36.00%	44.00%	40.00%
C12: 제품A 평균값	3.425542398	189.8267853	186.40	\$9.00	\$11.00	\$10.00
C13: 제품B평균값	16.70663096	176.5456968	159.84	\$11.03	\$13.48	\$12.25
C15: 제품A 수량	23.1774976	170.0748301	146.90	45.00	55.00	50.00
C16: 제품B 수량	30.5329996	162.7193281	132.19	31.50	38.50	35.00
C14: 제품C평균값	40.14658725	153.1057405	112.96	\$13.64	\$16.67	\$15.15
C17: 제품C 수량	48.04736933	145.2049584	97.16	18.00	22.00	20.00
C5: 조정된 시장리스크 할인율	138.2391267	57.029841	81.21	13.50%	16.50%	15.00%
C8: 금액약화율	116.803809	76.64095245	40.16	4.50%	5.50%	5.00%
C7: 연간 매출증가율	90.58835433	102.6854117	12.10	1.80%	2.20%	2.00%
C24: 감가상각비	95.08417251	98.1681552	3.08	\$9.00	\$11.00	\$10.00
C25: 상각	96.16356645	97.08876126	0.93	\$2.70	\$3.30	\$3.00
C27: 미자지급	97.08876126	96.16356645	0.93	\$1.80	\$2.20	\$2.00

Figure 5.4 – 감도표

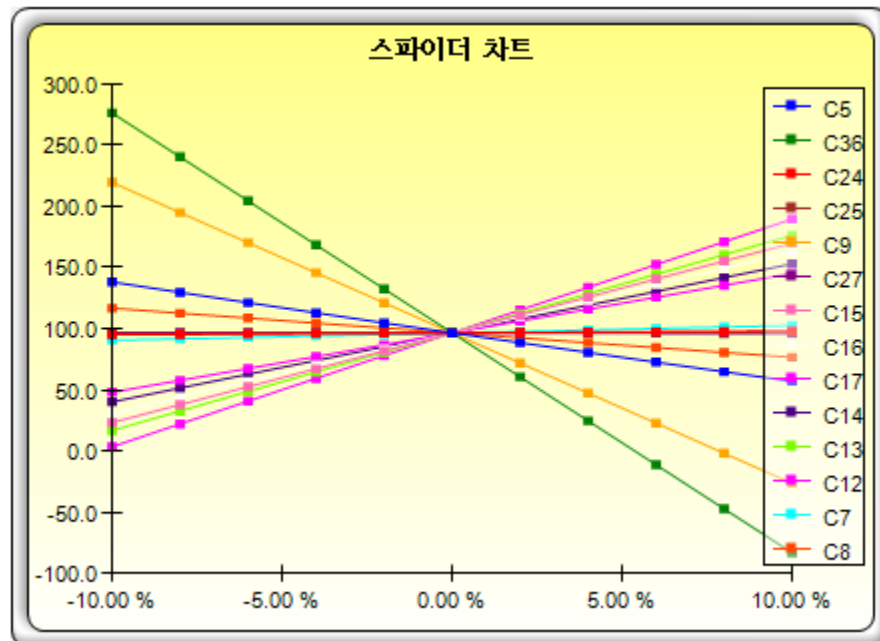


Figure 5.5 – 스파이더 차트

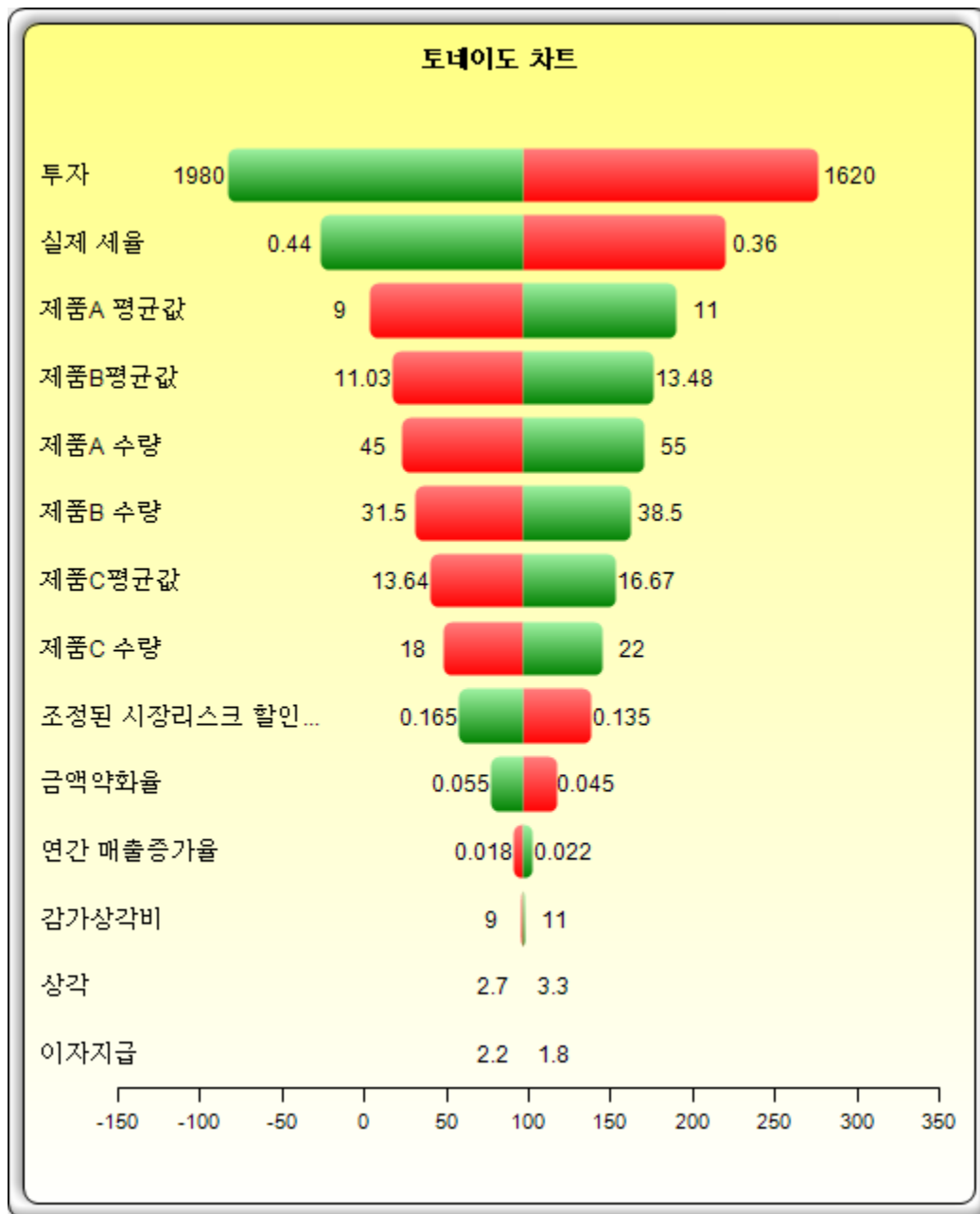


Figure 5.6 – 토네이도 차트

토네이도 차트는 독해하는 것이 간단하고, 스파이더 차트에서는, 모델에 비선형이 있는가 어떤가가 중요하게 됩니다. 예를 들어, Figure 5.7 는, 비선형이 꽤 명백 (그래프에서 선은 직선이 아닌 곡선입니다) 한 스파이더 차트를 표시합니다. 사용된 모델은, **토네이도와 감도 차트 (비선형)** 으로, Black-Scholes 옵션의 가격 모델을 예증 모델로서 사용하고 있습니다. 비선형은, 토네이도에서 확인되어, 모델에 있어서 중요한 정보가 되나, 모델의 활동 중에서 중요한 판단 메카를 부여합니다.

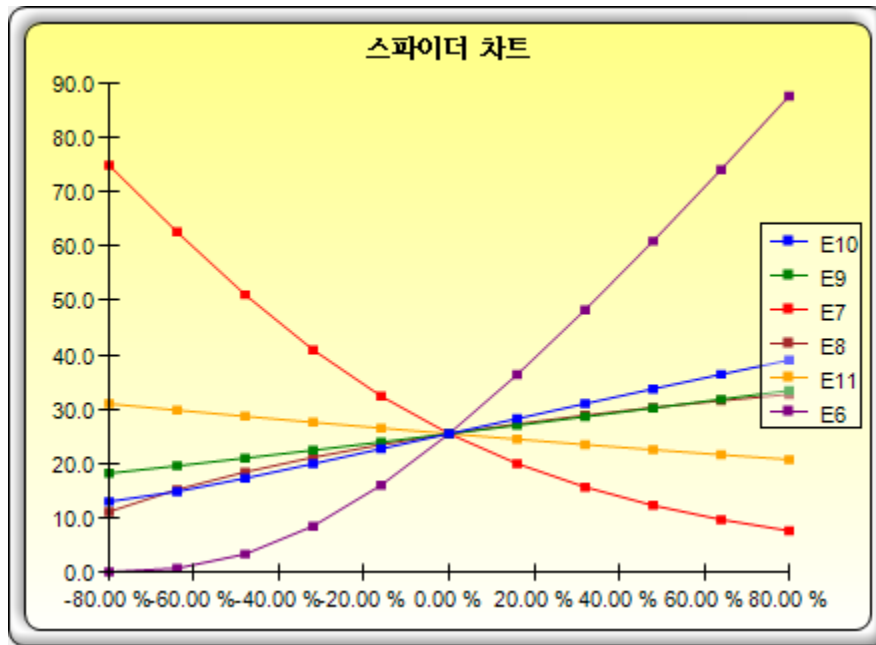


Figure 5.7 – 비선형적 스파이더 차트

이론:

관련하는 방법은 감도 분석입니다. 토네이도 분석(토네이도 차트와 스파이더 차트)가 시뮬레이션의 실행 전에 정적인 교란을 적용하는가 한편, 감도 분석은, 시뮬레이션의 실행 후에 작성된 다이나믹 한 교란을 적용합니다. 토네이도와 스파이더 차트는, 정적 교란의 결과로서, 각 선례, 또는 가정변수는 초기의 양을 한번 교란 시켜, 변동의 결과는 표에 표시됩니다. 한편, 감도 차트는 다이나믹한 교란의 결과로서, 복수의 가정은 동시에 교란되고, 모델과 변수의 사이에 상관으로 이것의 반복은, 결과의 변동에서 취득 될 수 있는 것을 의미하고 있습니다. 따라서 토네이도는, 시뮬레이션에 있어서 어느 변수가 가장 적절한 결과로 이끌어 주는가를 보여 주고 있습니다. 감도 차트는 결과에의 영향을 보여주는 한편, 상호하고 있는 복수의 변수는, 모델에 동시에 시뮬레이션됩니다. 이 효과는 명확히 Figure 5.8 에 표시되고 있습니다. 크리티컬한 성공의 드라이버의 순위표는, 전에 기술된 토네이도 차트의 예증에 유이한 것에 주목하여 주십시오. 단, 가정간의 상이에 상관이 추가되는 경우, Figure 5.9 와 같이, 매우 다른 결과를 표시하고 있습니다. 예증에 주목하여 주십시오. NPV 상에서 가격의 감소는 작은 영향을 불러 들입니다만, 입력 가정의 어느 것인가가 상관하는 경우에, 이것의 상관되는 변수 사이에 존재하는 반복은, 가격의 감소가 더욱 큰 영향을 불러 일으키도록 하여 줍니다.

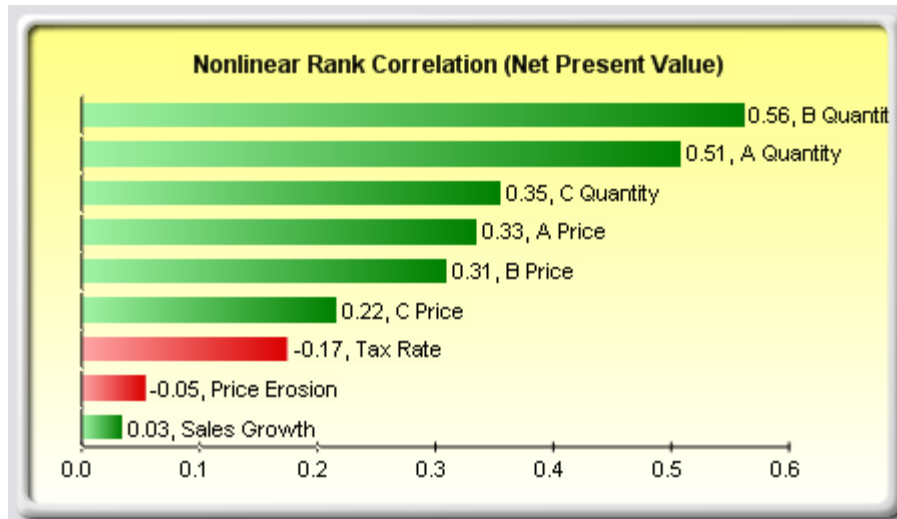


Figure 5.8 – 상관없는 감도 차트

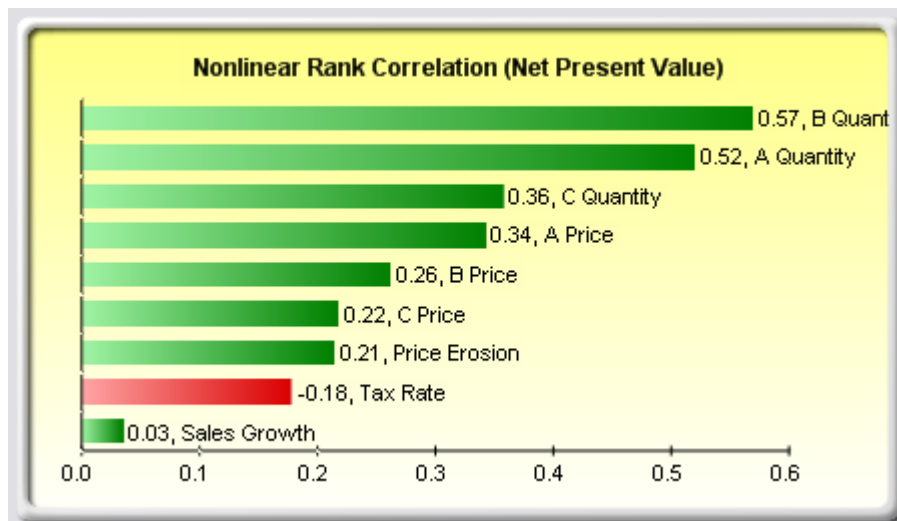


Figure 5.9 –상관과 감도 차트

순서:

- ❏ 모델을 작성하거나 열어 주십시오. 가정과 예측을 정하여, 시뮬레이션을 실행하여 주십시오 (여기에서의 예증은, 토네이도와 감도 차트 (선형) 파일을 사용하고 있습니다).
- ❏ **리스크 시뮬레이터 (Risk Simulator) / 툴 (Tools) / 감도 분석 (Sensitivity Analysis)** 을 선택하여 주십시오.
- ❏ 분석을 위하여 예측을 선택하고, OK (Figure 5.10)를 클릭하여 주십시오.

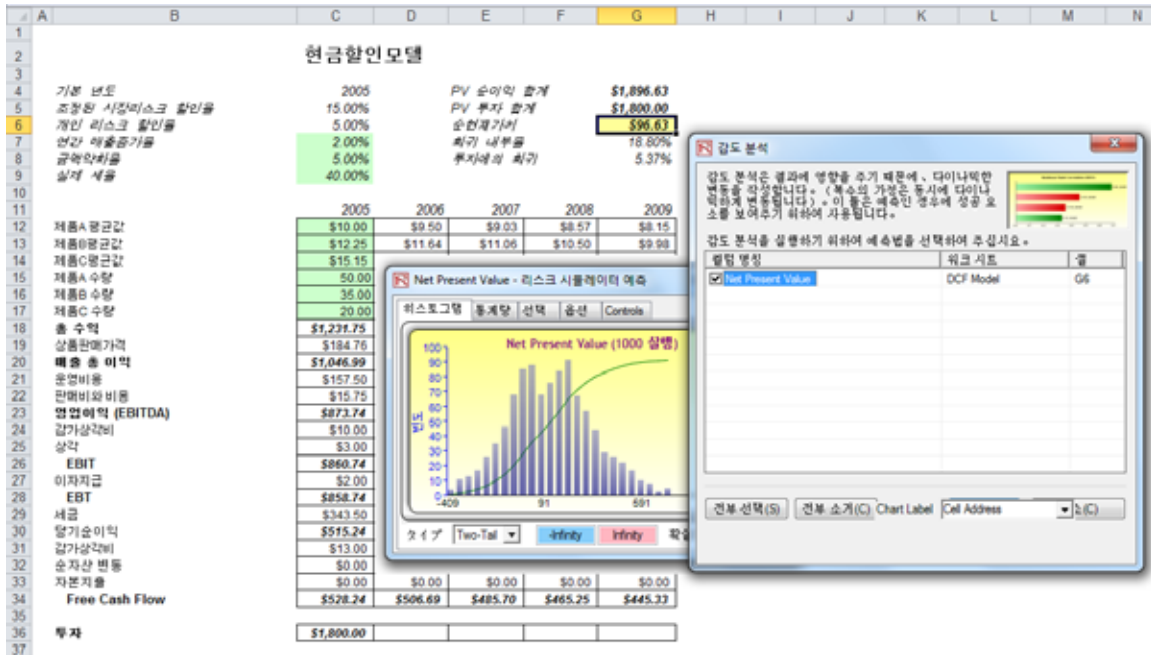


Figure 5.10 - 감도 분석의 실행

결과의 해석:

감도 분석의 결과는 레포트와 두 개의 키 차트를 구성합니다. 최초는, 가정 예측의 상관은 페어를 높은 것으로부터 낮은 것에의, 비선형 순위표 상관 차트 (Figure 5.11)를 구성합니다. 이 상관은 비선형 또한 논파라메트릭으로, 어느 분포 요건에서도 자유롭습니다(예, 와이블 분포와 가정은 베타 분포에서 비교할 수 있습니다). 이 차트에서의 결과는 전에 기술되는 토네이도 분석에 명백하게 유리합니다만 (물론, 알려져 있는 값이라면 초기에 정해져 있어서, 시뮬레이션되지 않은 투자 자본을 제거), 예외가 하나 있습니다. 감도 차트(Figure 5.11)에서는, 세율은 토네이도 차트(Figure 5.6)에 비하여, 무엇보다도 낮은 위치에 이행됩니다. 이것은, 세율이 유의한 영향을 불러 일으키지만, 일단, 모델에서의 다른 변수가 상호작용한 후, 세율의 영향을 감소시키는 것부터입니다 (이것은, 세율은, 이력적 세율의 경향에는 별로 변동이 없는 것처럼, 분포가 작고, 혹은, 세율은, 세금전의 수입의 백분위수의 값이 스트레이트로, 다른 선례의 변수 쪽이 커다란 효과를 가지고 있기 때문입니다). 이 예증은 시뮬레이션 실행후의 감도 분석의 적용은, 모델내에서 어떠한 상호 관계가 있는지, 그리고 특정의 변수의 영향이 유지 되는지를 확인하는 것으로 중요하다는 것을 증명합니다. 두 번째의 차트 (Figure 5.12)는, 설명된 백분위 수의 변동을 표시하고 있습니다. 즉, 예측의 변동 내에서, 어느 정도의 변동이, 변수 사이의 모든 상호 작용을 고려한 각 가정에 의하여 설명되는가를 표시하고 있습니다. 설명된 변동의 덧셈은 항상 100% (때에 따라, 모델에 영향을 주는 다른 요소가 있습니다. 직접 여기에서는, 채취할 수 없습니다)에 가까운 것에 주목하여 주십시오. 그리고 상관성이 존재하는가

어떤가, 그리고 이것의 덧셈은, 때에 따라서 100% (축척은 반복 효과의 실행을 위함)를 넘습니다.

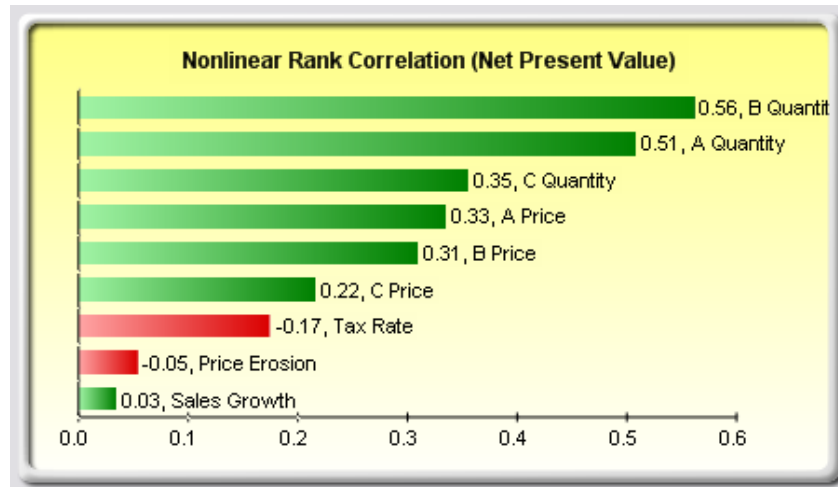


Figure 5.11 - 상호 순위 차트

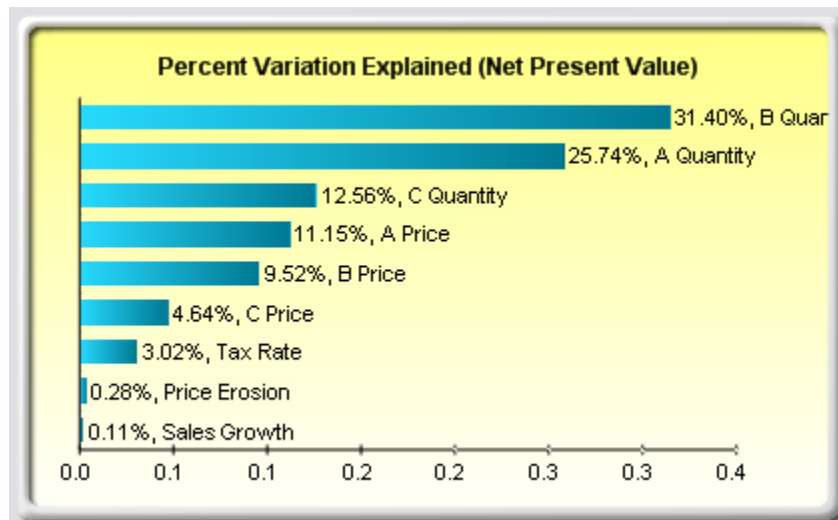


Figure 5.12 - 분산 차트에서의 공헌

메모:

토네이도 분석은, 시뮬레이션의 실행전에 수행되어, 감도 차트는 시뮬레이션의 실행후에 수행됩니다. 토네이도 분석에서의 스파이더차트는 비선형을 고려하는 것이 가능, 감도 차트의 랭크 상관관계가 비선형과 분포적으로 자유로운 조건이 고려됩니다.

:

이론:

다른 강력한 시뮬레이션 툴은 분포적 적합입니다. 즉, 어떠한 분포가, 모델의 특징의 입력 변수의 분포의 적용에 해당하겠습니까? 어느 것이 중요한 분포 파라미터입니까?

이력 데이터가 존재하지 않는 경우, 분석가는 의문이 되는 변수의 가정을 작성하지 않으면 안됩니다. 하나의 방법으로서, 전문가의 그룹이 각 변수의 행동의 추정이 맡겨진 Delphi 방법을 사용하는 것입니다. 예를 들어, 기계 공학자의 그룹이 검토, 및 엄격한 실험을 통하여 스프링 코일의 직경의 극단적 확률의 평가를 하는 것등. 이것의 값은 변수의 입력 파라미터 (예, 정규 분포와 0.5 과 1.2 의 사이의 극단값) 으로서 사용할 수 있습니다. 검증이 가능하지 않으면 (예, 마켓 셰어와 수입의 성장률) , 가능성이 있는 결과 추정의 관리에 의하여 작성할 수 있고, 최적의 케이스, 무엇보다 대략의 케이스, 그리고 나쁜 케이스의 시나리오를 부여합니다.

단, 신뢰 할 수 있는 이력 데이터가 존재하는 때에는, 분포적 적합을 행하는 것이 가능합니다. 이력 패턴과 이력 경향 자신의 반복은, 최적의 적합 분포와, 시뮬레이션 되는 변수의 최적의 정의의 중요 파라미터의 검출에 사용됩니다. Figures 5.13 에서 5.15 는, 분포 적합의 예증을 표시하고 있습니다. 이 일러스트는, 예증 파일의 **데이터의 적합** 파일을 사용하고 있습니다.

순서:

- 분포 시트를 열어, 적합한 데이터를 불러 들여 주십시오.
- 적합한 데이터를 선택하여 주십시오 (데이터는 하나의 종례의 복수 셀에 기입하지 않으면 안됩니다.)
- **리스크 시뮬레이터 (Risk Simulator) / 툴 (Tools) / 분포 적합(단일 변수) (Distributional Fitting (Single-Variable))** 를 선택하여 주십시오.
- 적합한, 어떤 특정의 분포를 선택하거나, 모든 분포가 선택되어 있는 디폴트를 유지하고, OK (Figure 5.13)를 클릭하여 주십시오.
- 적합의 결과를 확인하고, 희망하는 중요한 분포를 선택하여, OK (Figure 5.14)를 클릭하여 주십시오.

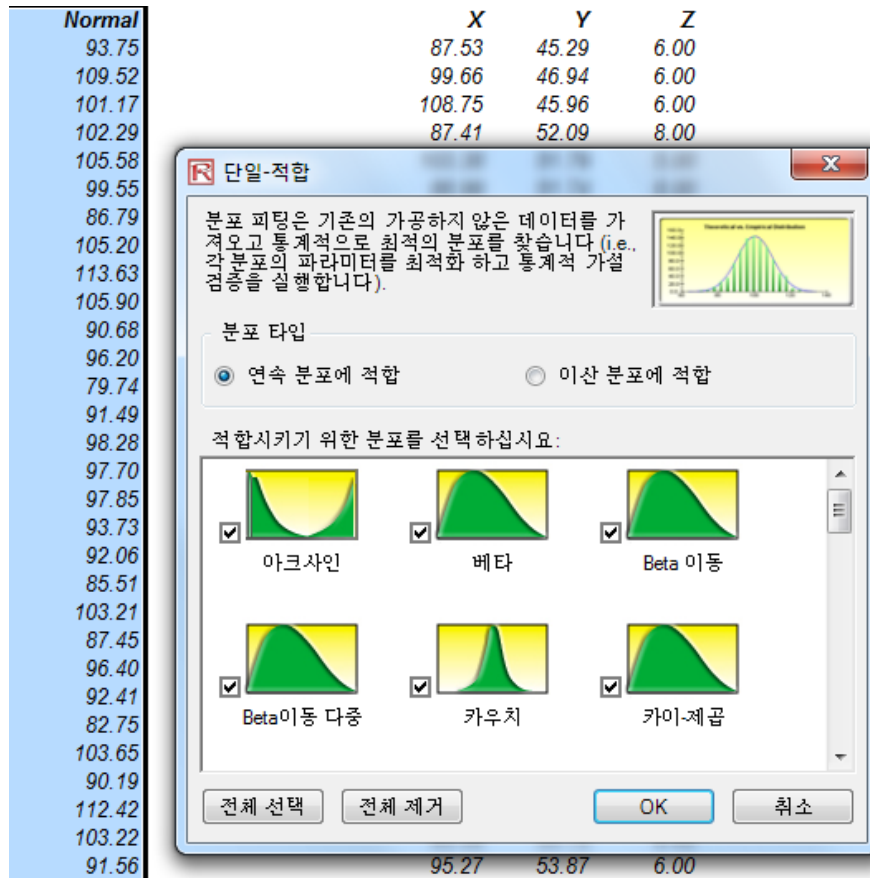


Figure 5.13 – 단일 변수의 분포 적합

결과의 해석:

무귀 가설은, 적합 대상의 분포가, 적합 대상의 샘플 데이터에서 끌어오는 모집합과 같은 분포를 가지고 있는가 어떤가를 검정할 수 있습니다. 따라서, 계산된 p-값이 크리티컬한 알파 레벨(일반적으로 0.10 이나 0.05)보다 적은 경우는, 분포는 틀린 것을 표시하고 있습니다. 한편, p-값이 높으면 무엇보다 좋은 분석은 데이터에 적합합니다. 대부분, p-값을 설명하는 백분위로서 생각할 수 있는 것이 가능, p-값이 0.9727 (Figure 5.14) 인 경우, 99.28 의 평균값과 10.17 의 표준 편차를 가진 정규 분포가, 데이터의 변동의 97.27%를 설명하고 있어, 특별히 좋은 적합을 표시하고 있습니다. 어느 쪽의 결과(Figure 5.14)도, 레포트 (Figure 5.15)도 검정 통계, p-값, 이론적 통계 (선택된 분포에 기초하여), 경험적인 통계 미가동 데이터에 기초하여, 초기 데이터(사용되는 데이터의 기록의 유지)와 완성된 가정과 중요한 분포의 파라미터(예, 자동적인 가정의 생성의 선택과 시뮬레이션 프로필이 이미 존재하는 경우)를 표시하고 있습니다. 또한, 결과는 선택된 모든 분포와 어느 정도의 데이터에 적절히 적합하는가를 순위표에 표시합니다.

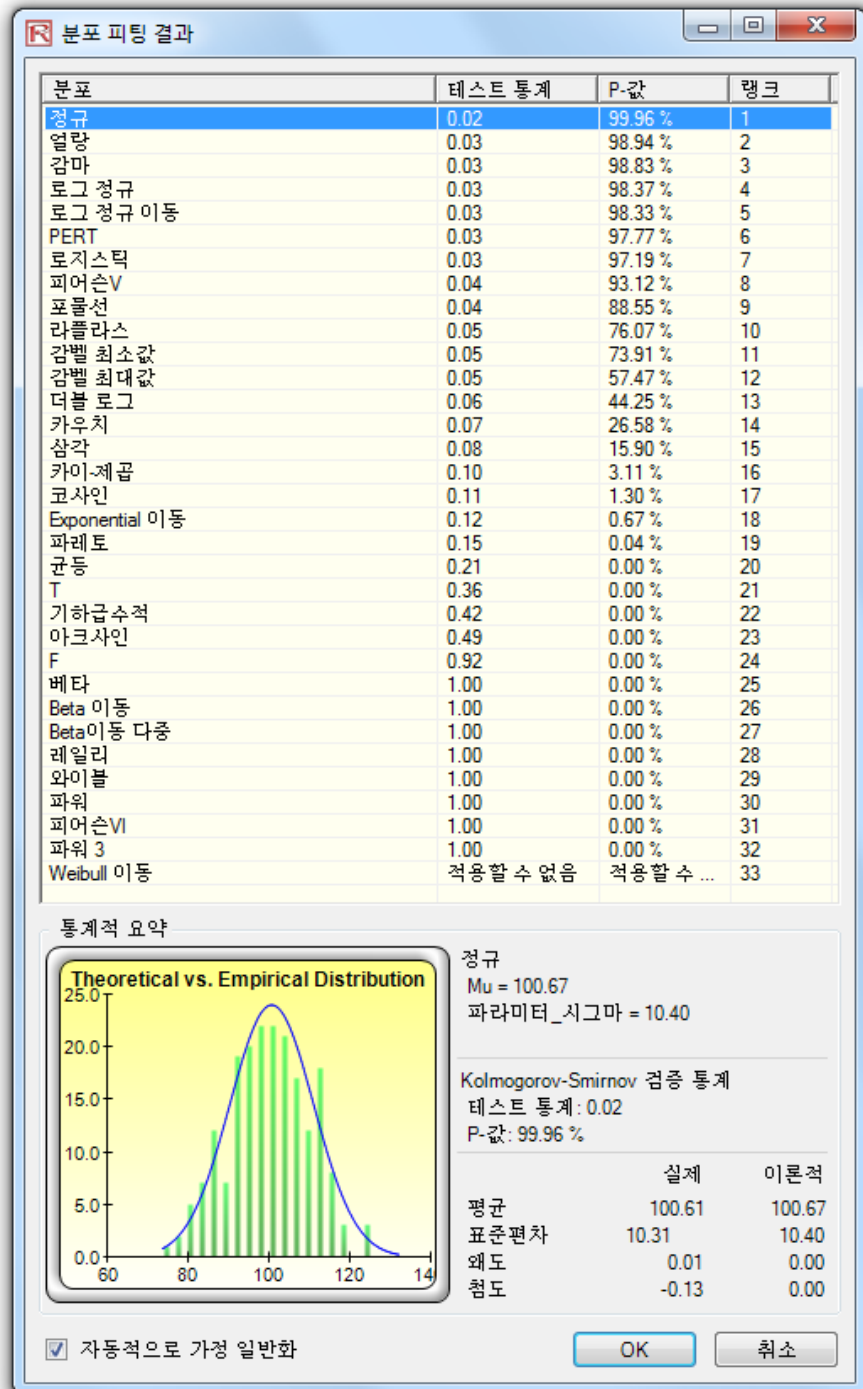


Figure 5.14: 분포 적합의 결과

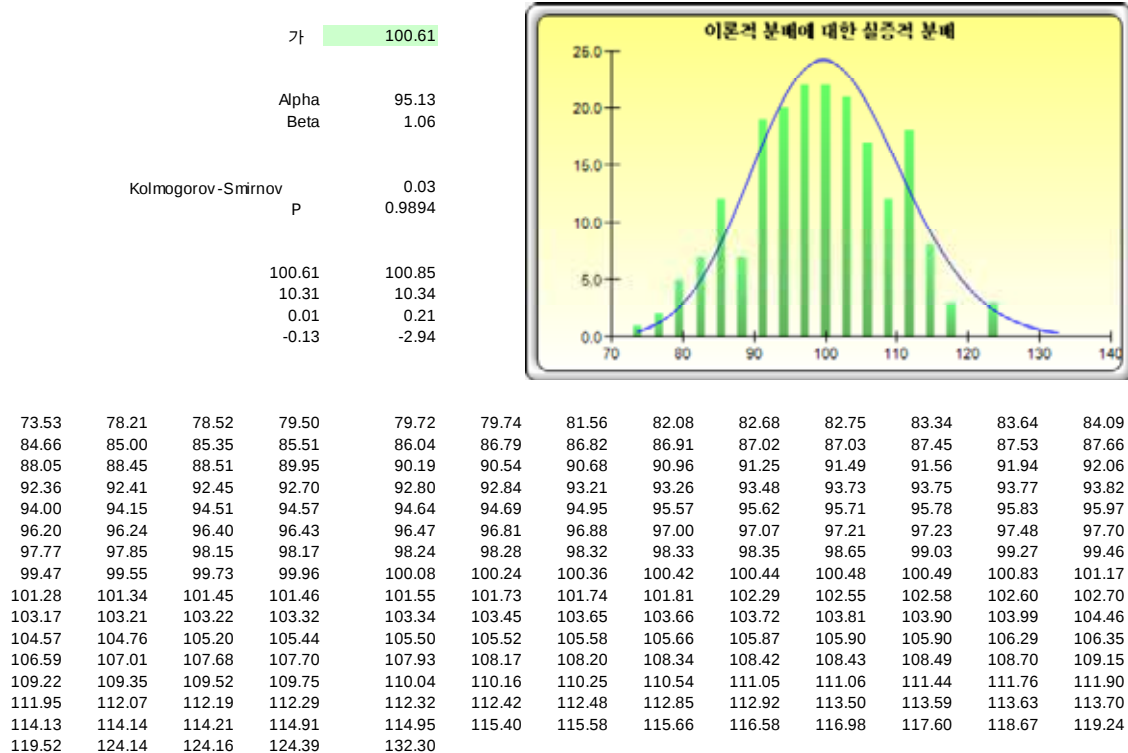


Figure 5.15: 분포 적합의 레포트

복수의 변수의 적합 과정은, 단일 변수가 적합한 때와 같이 꽤 비슷합니다. 어차피, 데이터는 열로 편집되지 않으면 읽힌다고 (예 : 각 변수는 하트의 열이 된다) 하면, 다른 변수는 한번에 모두 적합하게 됩니다.

순서:

- 분산 시트를 열어, 적합한 데이터를 불러 들여 주십시오.
- 적합한 데이터를 선택하여 주십시오 (데이터는 복수의 종열과 횡열에 들어가지 않으면 읽힙니다.)
- **리스크·시뮬레이터 (Risk Simulator) / 툴 (Tools) / 분포 적합 (복수의 변수) (Distributional Fitting (Multiple Variables))** 를 선택하여 주십시오.
- 데이터를 재검토하여, 관련 있는 분포 타입을 선택하고, OK 를 클릭하여 주십시오.

메모:

분포적 적합 루틴에 이용되고 있는 통계 순위 방법은 제곱 검정과 Kolmogorov-Smirnov 검정을 사용합니다. 전자는 이산계 분포를 후자는 연속 분포를 검정하기 위하여 사용됩니다. 간결하게 말하면, 내부 최적화 루틴과 맞춘 가설 검정은 검정된 각

분포의 파라미터를 보여 주기 위하여 사용됩니다. 그리고, 결과는 순위표에 기입됩니다.

Boot Strap

이론:

Boot Strap 시뮬레이션은 예측 통계의 정도, 신뢰성 및 샘플 가공 데이터등을 추정하는 간단한 테크닉법입니다. 기본적으로, Boot Strap 은 가설 검정으로 사용됩니다.

고전적인 방법론에서는 샘플 통계의 신뢰성을 기술하기 위하여 수학적 방식을 사용하고있습니다. 이 방법은 샘플 통계의 분포가 정규 분포에의 통계의 표준 에러, 혹은 신뢰 구간의 계산에 비교적 쉽게 가까이 갈 수 있습니다. 어쨌든, 통계의 샘플 분포가 정규 분포가 아니고, 또는 간단히 찾을 수 없는 경우 이 고전적 방법은 무효이거나, 사용이 곤란하게 되어 있습니다. 역으로, BootStrap 이 샘플 통계 데이터를 몇번이고 반복하여 샘플로 하고, 그러한 샘플을 기초로 여러 가지 통계 분포를 만들어 내어, 실천적으로 분석합니다.

순서:

- 시뮬레이션을 실행하여 주십시오.
- *리스크·시뮬레이터 (Risk Simulator) /툴 (Tools) /논 파라메트릭 Boot Strap (Nonparametric Bootstrap)* 을 선택하여 주십시오.
- Boot Strap 의 실행을 위하여 한 개만의 예측을, 통계를 선택하고, Boot Strap 시행 수를 기입하고 OK 를 클릭하여 주십시오. (참조 Figure 5.16)

Revenue	\$200.00	Revenue	\$200.00
Cost	\$100.00	Cost	\$100.00
Income	\$100.00	Income	\$100.00

이 모델을 복사하기 위하여, 시뮬레이션 프로필을 생성하여 시작하고 명칭을 만듭니다.
(Simulation I New Profile), 그리고, 평균 종자값을 123456 로 세팅합니다. 다음,
revenue 셀을 선택하고 평균 200, 표준편차 20 의 정상 분포를 제공합니다.
(Revenue 셀을 선택하고 Simulation I Set Assumption 을 클릭한 후,
Normal 을 선택 하고 관련 파라미터를 입력합니다). 그리고, 각 코스트 셀을 Normal assumption 으로 정의합니다.
결과적으로, 두 가지 수입 셀에 예측 값을 정의하고, 시뮬레이션을 실행합니다.

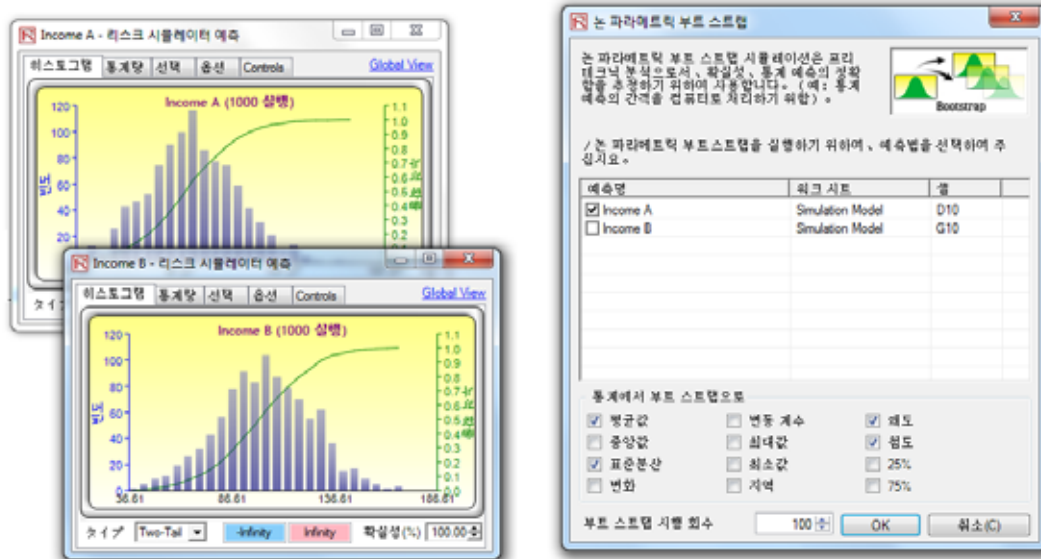


Figure 5.16 – 논파라메트릭 Boot Strap 의 결과

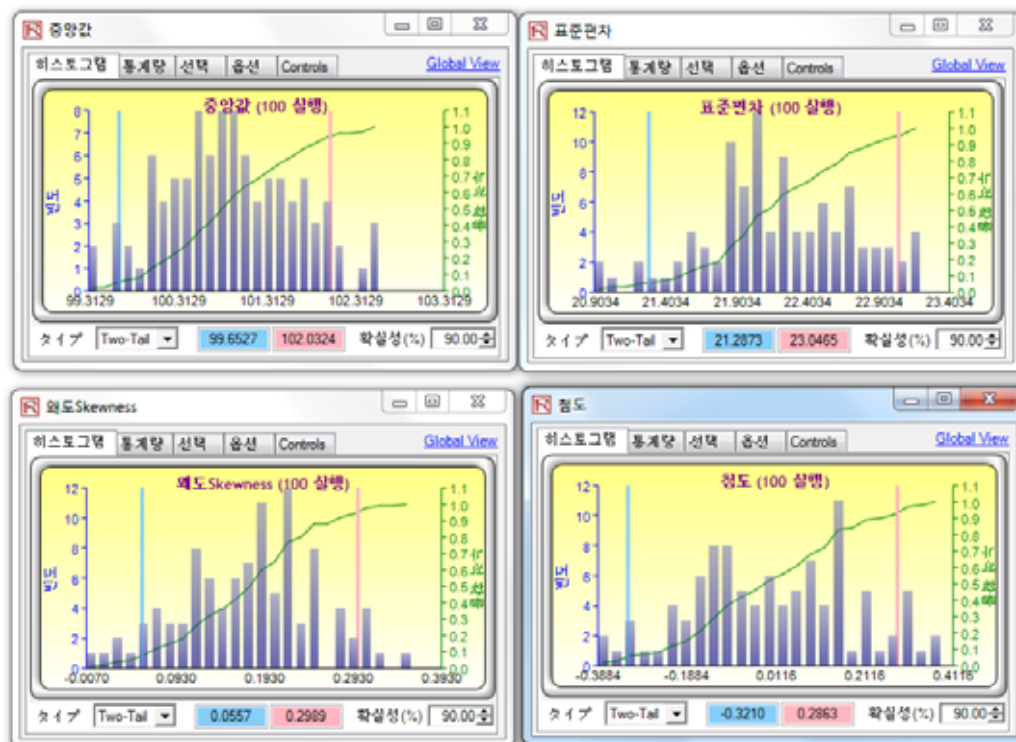


Figure 5.17 –Boot Strap 시뮬레이션의 결과

결과의 해석:

기본적으로, 논파라메트릭 Boot Strap 은 시뮬레이션에 기초한 시뮬레이션으로 불리워집니다. 시뮬레이션의 실행 후, 결과에서 일어나는 통계가 표시되어 있습니다. 그러나, 각 통계의 정도와 통계적 의미는 때대로 해결되지 않는대로인 경우가 있습니다. 예를 들어, 시뮬레이션 실행의 왜도의 통계가 -0.10 인 경우, 정말로 분포는 마이너스의 기울기입니까, 또는 최소한의 마이너스 수는 랜덤 시행에 기인하는 것입니까? 즉, 통계가 -0.15, -0.20 의 경우에는 어떻게 됩니까? 즉, 어느 정도 충분히 거리가 있으면, 분포에 마이너스의 경사가 있다고 생각됩니까? 이 의문은 통계에도 적용할 수 있습니다. 하나의 분포는 통계적으로 다른 분포와 같다고 하면 통계적 관계가 계산될 수 있을까, 또한 이것은 전혀 상위적입니까? Figure 5.17 은 Bootstrap 의 샘플을 참조하고 있습니다. 또한, 90%의 통계의 왜도의 신뢰는 -0.0189 에서 0.0952 의 사이에 있어, 0 값도 신뢰도안에 들어갑니다. 즉, 90%의 신뢰는 통계의 왜도에 통계적으로 제로와 다르지 않은 것을 표시하고 있고, 이 분포는 대칭적으로 기울기가 없는 것으로 생각해도 괜찮습니다. 역으로 값이 0 을 넘는 경우는, 그 반대를 표시하고, 분포는 비대칭적으로 비뚤어진 것을 표시합니다. (예측 통계가 플러스인 경우 플러스는 없고, 예측 통계가 마이너스인 경우 기울기도 마이너스인 경우가 있습니다) .

메모 :

Bootstrap 의 유래는 “자신의 힘으로 자기 자신을 끌어 올린다”라는 의미가 있고, 이 방법은, 통계학의 정도를 분석하는데 통계자체의 분포를 사용합니다. 논파라메트릭 시뮬레이션은 각 골프 볼이 이력 데이터 포인트에 기초하여 있다고 하고, 교환으로 큰 상자에서 골프 볼을 단순히 랜덤 방식으로 선택하는 것입니다. 예를 들어, 상자안에 365 개의 골프 볼이 들어가 (365 의 이력 데이터 포인트를 표시하고 있습니다) 있다고 상상해 봅시다, 랜덤 방식으로 선택 된 각 골프 볼의 값이 화이트 보드에 써져 있다고 합시다. 교환되지 않은 선택된 골프 볼의 결과는 보드의 365 수의 열의 일렬에 기입하고, 중요한 통계가 이것의 365 열상에 (예 : 평균값, 중간값, 표준 편차 등) 계산됩니다. 그리고, 예를 들어 과정이 5,000 회 반복된다고 합시다. 그렇게 되면 화이트 보드는 365 열과 5,000 의 종렬에 기입하는 것이 됩니다. 따라서, 통계의 결과는 (이것은 5,000 의 평균값, 5,000 의 중간값, 5,000 의 표준 편차등이 있는 것을 인식하고) 표에 기입되어, 각각의 분포는 표시되어 있습니다. 중요한 통계는 이후에, 표에 표시됩니다. 이 결과에 의하여, 시뮬레이션의 예측의 신뢰도를 확인할 수 있습니다. 즉, 10,000 회 시행의 시뮬레이션에서는, 결과로서 일어나 예측 평균이 5.00 달러라는 것을 전제로 하면, 분석자는 어느 정도 결과에 대하여 정밀도를 가지고 있겠습니까? Boot Strap 은 계산된 통계값의 분포의 표시도 포함하고 있고, 신뢰 구간을 확인시켜 주십시오. 최후에, Boot Strap 의 결과는, 통계의 대수의 법칙과 중심 극한 정리에 의하면, 샘플 평균값은 샘플 사이즈가 늘어나면, 불편 추정률로, 진짜 모집단 평균값에 거의 가깝고, 비슷한 것을 알수가 있습니다.

가

이론:

가설 검정은 두 개의 분포의 평균값 및 분산의 검정을 행할 때에는, 그것이 통계적으로 같은가, 다른가를 표시하여 줍니다. 전혀 다른 두 개의 예측의 평균값 및 분산의 차이는 랜덤 시행이나, 실제로 통계적으로 전혀 다른 값에 기초하고 있다는 것을 생각할 수 있습니다.

순서:

- 시뮬레이션 실행
- 리스크·시뮬레이터 (Risk Simulator) / 툴 (Tools) / 가설 실험 (Hypothesis Testing) 을 선택
- 동시에 실험을 실행하는 두 개의 예측을 선택, 가설 실험의 타입을 선택하고, OK 를 클릭하여 주십시오(Figure 5.18)。

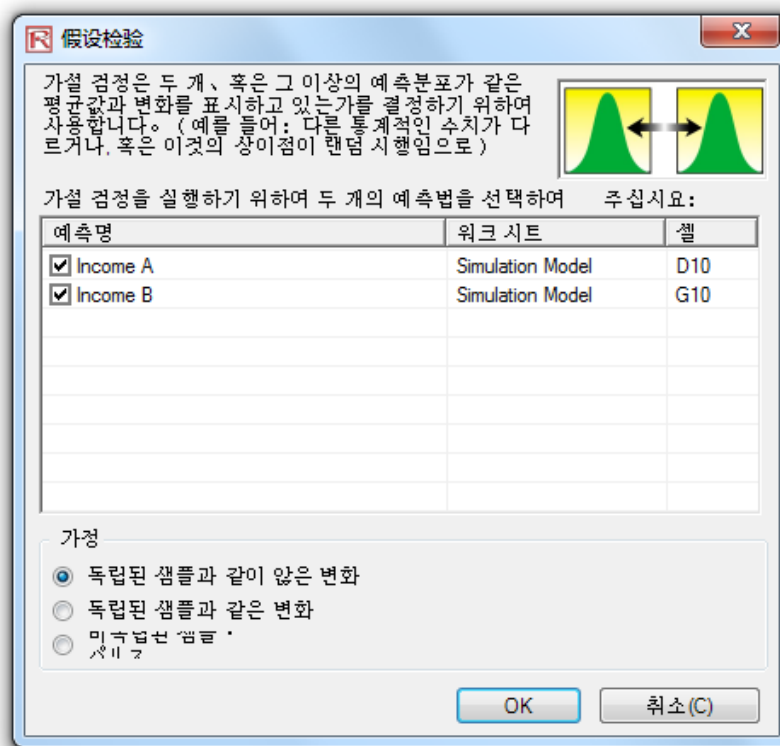


Figure 5.18 – 가설 실험

레포트 해석:

양측 가설의 검정은 귀무 가설상 실행되고 있습니다(H_0)。 이 두 개의 모집단(변수)는 통계적으로 비슷하다도 말하는 것을 의미하고 있고、대립 가설(H_a)은 두 개의 모집단(변수)는 통계적으로 비슷하지 않다는 것을 의미하고

리스크 시뮬레이터의 데이터 취득의 순서를 사용하면 간단히 시뮬레이션의 미가공 데이터의 취득이 가능합니다. 가정의 데이터도 예측의 데이터도 취득 가능합니다만、미리 시뮬레이터가 실행되어 있는 것을 확인하여 주십시오。취득된 데이터는 나중에 다양하게 다른 분석에서도 사용할 수 있습니다。

순서:

- 모형을 열거나 작성한 후、가정과 예측을 설정하고、시뮬레이션을 실행하여 주십시오。
- **리스크 · 시뮬레이터 (Risk Simulator) / 툴 (Tools) / 데이터 취득 (Data Extraction)** 을 선택하여 주십시오。
- 데이터를 취득하고자 하는 가정、예측、또는 양방을 선택하여、OK 를 클릭하여 주십시오。

취득한 데이터는 여러 가지 포맷으로 보존할 수 있습니다。

- 미가공 데이터에서 새로운 워크 시트에 표시되어, 시뮬레이터된 값 (가정 그리고 예측의 양방) 은、보존、또는 분석할 수 있습니다。
- 플랫 텍스트 파일 (flat text file) 。 데이터는 다른 분석 소프트웨어에 Export 됩니다。
- 리스크 · 시뮬레이터 파일。 **리스크 · 시뮬레이터 (Risk Simulator) / 툴 (Tools) / 데이터를 열고·인포트 (Data Open/Import)** 를 선택 (가정 혹은 예측) 하면 언제든지 결과를 회복하는 것이 가능합니다。

세번째의 옵션은、시뮬레이터된 데이터 결과를 *.risksim 파일로서 보존하는 것입니다。데이터 결과는 언제든지 회복하는 것이 가능하고、시뮬레이션을 매회 실행시킬 필요가 없습니다。Figure 5.21 의 다이알로그 · 박스에서는 데이터를 취득、익스포트、취득 데이터 결과의 보존의 화면이 표시됩니다。

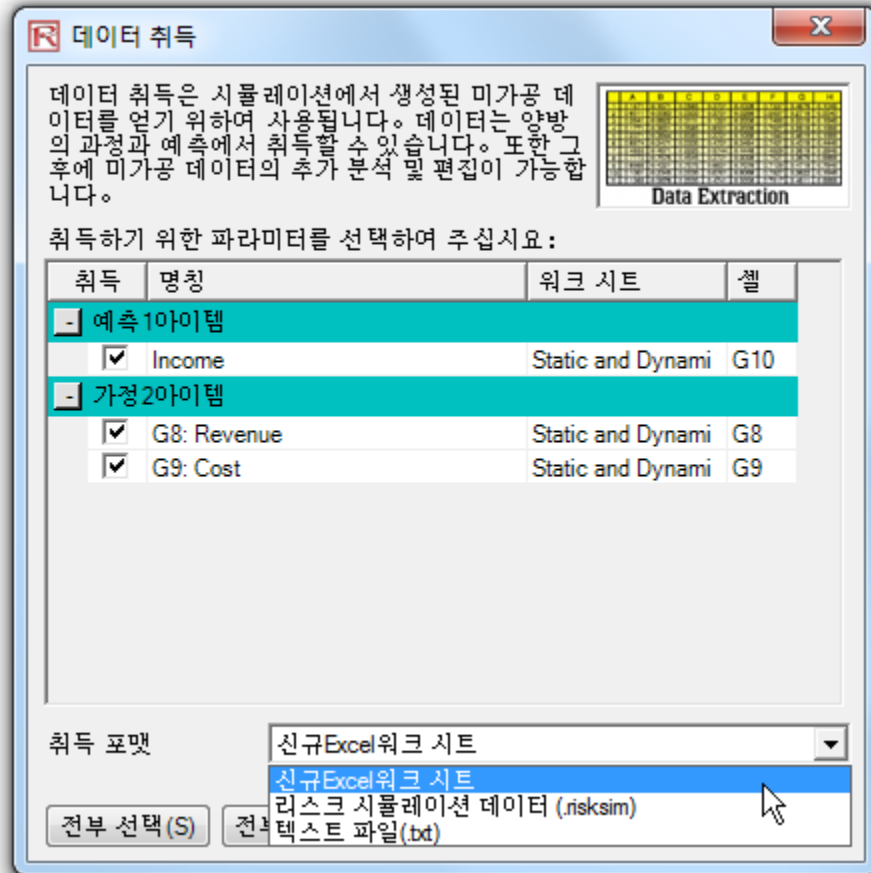


Figure 5.21 – 시뮬레이션 레포트의 샘플

시뮬레이션 실행 후, 설정된 가정, 예측, 시뮬레이션의 결과등의 레포트의 작성이 가능합니다.

레포트의 작성의 방법:

- ❏ 모형을 열거나 작성한 후, 가정과 예측을 설정하고, 시뮬레이션을 실행하여 주십시오.
- ❏ **리스크 시뮬레이터 (Risk Simulator) / 레포트의 작성 (Create Report)** 을 선택하여 주십시오

- DCF Model

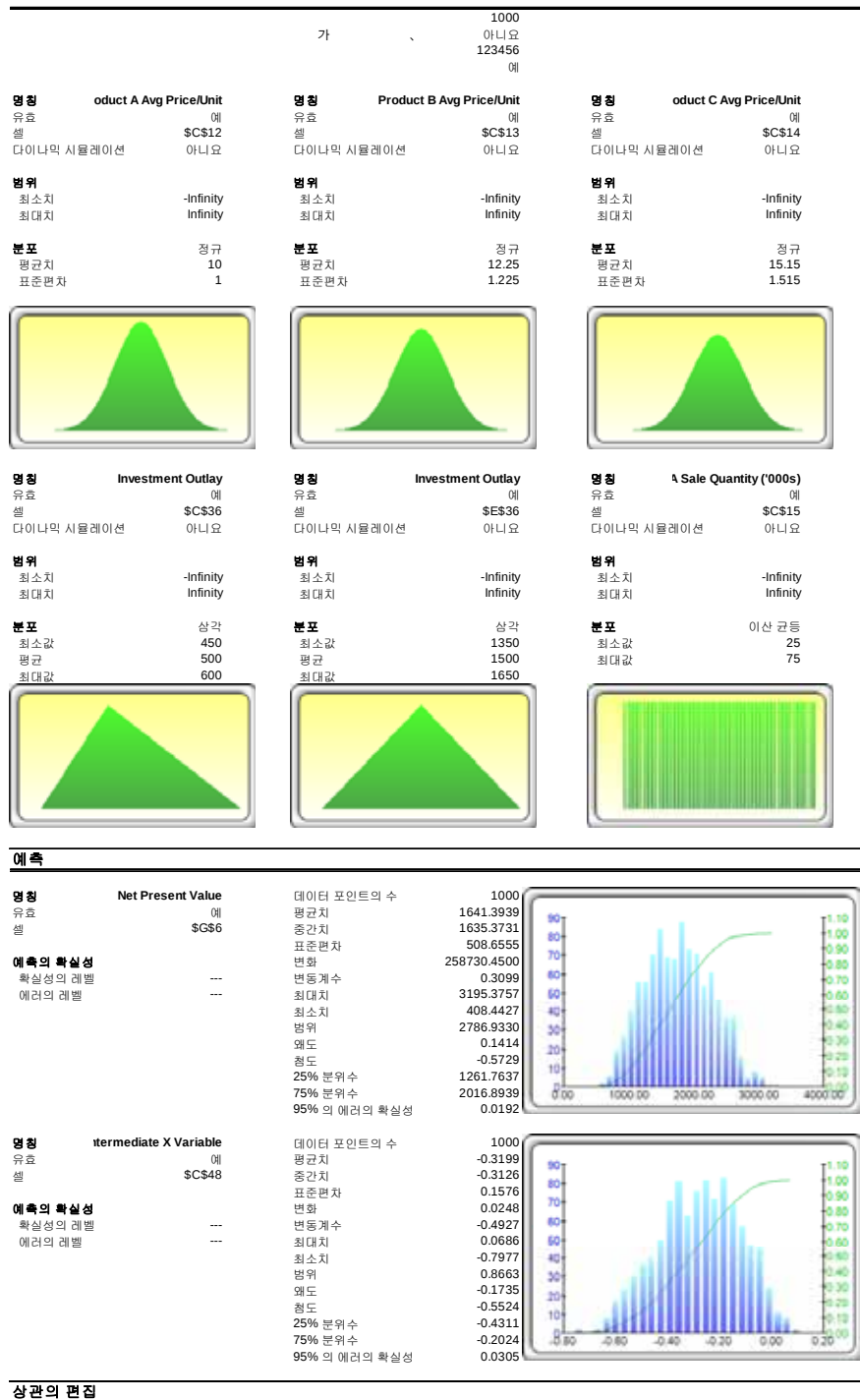


Figure 5.21 – 시뮬레이션 레포트의 샘플

리스크 시뮬레이터에서의 고도의 분석 툴은, 데이터의 계량 경제학의 성질을 정하기 위하여 사용됩니다. 진단은, 불균일 분산, 이상치, 지정된 에러, 소수, 계절성과 확률의 성질, 정상성, 에러의 구형과 다중 공선성을 포함하고 있습니다. 각 검정은, 그것의 모델의 레포트에서 자세히 기술되어 있습니다.

순서:

- 샘플 모델을 열고(리스크 시뮬레이터 (Risk Simulator) | 샘플 (Examples) | 회귀 진단 (Regression Diagnostics))、시계열의 데이터 워크 시트로 변수 명칭을 포함한 데이터를 선택하여 주십시오(셀 C5:H55).
- 리스크 시뮬레이터 (Risk Simulator) | 툴 (Tools) |진단 툴 (Diagnostic Tool) 을 선택하여 주십시오.
- 데이터를 확인하고, 하기의 메뉴에서 종속 변수 Y 를 선택하고, 종료 후、OK 를 클릭하여 주십시오(Figure 5.22).

다중 회귀 분석 데이터 세트

독립변수 Y	변수 X1	변수 X2	변수 X3	변수 X4	변수 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				
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				

독립변수 Y	변수 X1	변수 X2	변수 X3	변수 X4	변수 X5
183	1.578	20.5	2.7		
417	1.202	10.9	5.5		
233	1.109	123.7	7.2		

Figure 5.22 – 에디터의 진단 툴의 실행

예측과 회귀 분석의 일반적 위반은 불균일 분산으로, 에러의 변동은 시간과 같이 증가하는 것을 표시합니다. (참조, Figure 5.23 진단 툴을 사용한 검정 결과의 표시). 한편, 종속의 데이터 분산의 폭은 증가, 및 시간과 같이

넘어버려, 일반적으로, 결정 계수 (R^2 - 평방된 계수) 는 불균일 분산이 존재하는 때에 유의하게 하강합니다. 종속 변수의 분산이 일정하지 않을 때, 에러의 분산도 일정하지 않습니다. 종속 변수의 불균일 분산이 강조하면 이것의 효과는 그렇게 심각하지 않고: 최고 이승 추정은 또 바이어스 되지 않고, 기울기와 절편의 추정치, 에러가 정규 분포로 되지 않는 경우는, 어느쪽도 정규 분포가 되고, 또는, 에러가 정규 분포가 되지 않는 경우, 자가는 정상 점근적 분포 (데이터 포인트의 수가 커지게 되는 것처럼) 가 보여집니다. 기울기와 전체적인 변수의 값의 분산 추정은, 불확실한 가능성이 있습니다만, 종속 변수의 값이 이제부터의 평균값에 대조적이라면, 불확실성은 별로 커다란 가능성은 없습니다.

만약 데이터 포인트가 작음(소수)경우는, 가정의 위반의 검출이 곤란하게 됩니다. 작은 샘플 사이즈에서는, 비정규성, 및 분산의 불균일 분산과 같이 가정의 위반이 존재한다고 해도 검출이 어려워 집니다. 작은 수의 데이터 포인트와 선형 회귀는 과정의 반복 위반에서 방어는 별로 강하지 않습니다. 작은 데이터 포인트라면, 데이터에 적합한 선이 매치하거나, 혹은, 어떤 비선형 공식이 무엇보다도 적절한가를 정하는 것이 어려워 집니다. 하나의 가정이라도 위반되지 않았다고 하면, 데이터 포인트의 소수상의 선형 회귀는, 기울기와 제로 사이의, 예를 들어, 기울기가 제로에 동일하지 않다고 하면, 차이를 구분할 수 있는 충분한 힘을 가지고 있지 않습니다. 이 힘은 남아 있는 에러, 독립 변수에서 관측된 분산, 검정의 알파 레벨의 선택된 유효성과 데이터 포인트의 수에 의하여 변합니다. 힘의 감소는, 잔여의 분산의 증가, 유의 레벨의 감소 (예, 검정이 무엇보다도 어려워지는 것과 같음) 와 같이 감소하여, 관측된 독립 변수의 증가에서의 분산과 데이터 포인트의 수의 증가와 같이 증가합니다.

값은, 이상치의 존재가 있기 위하여, 같이 분포하지 않을지도 모릅니다. 이상치란 데이터의 비정상 값을 표시합니다. 이상치는, 적합한 기울기와 절편상에 강력한 영향을 불러 일으키고, 데이터 포인트의 크기에 빈약한 적합을 부여합니다.

이상치는, 잔여의 분산의 추정을 증가하는 경향을 보여, 귀무 가설의 거절의 찬스를 낮게 합니다. 예를 들어, 예언 에러의 확률을 높게 하는 등. 에러의 기록은, 수정이 가능하지 않으면 안되고, 또한, 같은 모집단에서 샘플되지 않은 종속 변수의 값의 기록은 중요합니다. 일견, 이상치는, 같은 모집단이지만 비정상적 모집단에서의 종속 변수의 값에서 유래하는지도 모릅니다. 단, 포인트로서, 분산 플롯에서의 이상치의 필요하지 않은 독립, 및 종속 변수의 어느쪽도 희귀한 값이 됩니다. 회귀 분석에서는, 적합한 라인이 이상치에 대단히 민감하지 않으면 안됩니다. 즉, 적은 평방의 회귀는 이상치에 대하여 저항이 없고, 따라서, 어느 것도 적합한 기울기의 추정이 되지 않습니다. 다른 포인트에서 종속변 소거된 포인트는, 데이터의 잔여의 전체적인 선의 트렌드를 따르는 대신에도, 특히 포인트가 데이터의 중심에서 횡선적으로 떨어져 있는 경우, 적합한 선을 포인트에 가깝게 통하는 원인이 됩니다.

단, 이상치가 소거된 경우, 좋은 방법을 선택할 필요가 있습니다. 그러나, 대부분의 경우에 이상치가 소거되는 경우, 회귀의 결과는 향상됩니다만, 먼저 경험적인 설명이 존재하지 않으면 안됩니다. 예를 들어, 어떤 특정의 회사의 주식 리턴의 실행을 돌리면, 주식 시장에서의 하강에 의하여 이상치가 포함되어 있지 않으면 안됩니다. 이것은, 비즈니스 싸이클내에서 어찌할 수 없는 것과 같은 진짜 이상치는 아닙니다. 이러한 이상치를 맞추어, 회귀 방식을 사용하고 회사의 주식에 기초한 한 사람의 퇴직기금을 예측하는 것은, 최대의 부정확한 결과를 불러옵니다. 한편, 이상치는, 비정상적인 비즈니스 코디네이션 (예, 합병 및 획득) 에 의하여 생기고, 이 비즈니스의 구성의 변환은, 반쪽되지 않도록 예측되지 않기 때문에, 이것의 이상치는 소거하는 것이 가능, 최고에 데이터를 청결히 해서 회귀 분석을 실행합니다. 여기에서의 분석은, 이상치가 검출되지 않고, 사용자의 취향에 의하여 이것을 남길 것인가 소거할 것인가를 선택할 수 있습니다.

때대로, 종속, 그리고 독립 변수의 비선형적인 관계는, 비선형적인 관계보다도 적절합니다. 어떤 케이스로 하여도, 선형적인 회귀를 실행하는 것이 최적은 아닙니다. 만약 선형 모델이 바르지 않다면, 기울기와 절편의 추정과 선형적인 회귀에서 적합한 값은 바이어스되어, 적합한 기울기와 절편의 추정은 유효하지 않습니다. 독립, 또는 종속 변수에 한정된 범위를 넘고, 비선형 모델은, 선형 모델에 가까우나 실제로 비선형 기본), 정밀가 있는 예고를 위하여는, 데이터에 있어서 적절한 모델의 선택이 필요합니다. 회귀의 실행 전에, 데이터에 비선형의 변환이 먼저 적용되지 않으면 안됩니다. 독립 변수 (다른 방법은, 평방근, 및 독립 변수를 두번째, 또는 세번째의 힘으로 상승합니다) 의 자연 대수를 얻는 것은, 회귀, 또는 예측을 비선형적으로 변환시킨 데이터를 사용하여 실행하는 것입니다.

Y	가		가	
	W-	p-value	p-value	
X1	0.2543	-7.86	671.70	2
X2	0.3371	-21377.95	64713.03	3
X3	0.3649	77.47	445.93	2
X4	0.3066	-5.77	15.69	3
X5	0.2495	-295.96	628.21	4
		3.35	9.38	3
				0.2458
				0.0335
				0.0305
				0.9298
				0.2727

Figure 5.23 – 이상치의 검정 결과, 불균일 분산, 소수와 비선형

시계열 데이터를 예측하기 위하여 다른 일반적인 방법은, 독립 변수의 값이 정말로 각자 독립하고 있는가, 또는 종속하고 있는가 어떤가를 확인하는 것입니다. 시계열을 통하여 취득한 종속 변수의 값은 자기 상관되지 않으면 안됩니다. 연속적으로 상관된 종속 변수의 값, 기울기와 절편의 추정은 언바이어스되지만, 이것의 예측의 추정과 분산은 불확실함으로, 특정의 통계적인 최선으로 적합한 검정의 평가는 다릅니다. 예를 들어, 이율, 인플레이션의 비율, 판매, 수입과 시계열 데이터에 관련하는 다른 비율 등, 현재의 주기의 값은 선주기의 값에 관련 되어 (분명히 3 월의

인플레이션은 2 월의 레벨과 관련되어, 이것 자체가 1 월의 레벨에 관련 되는 등) 、일반적으로 자기 상관됩니다. 이 관계를 무시하기 위하여, 바이어스를 그리고 정밀도가 적은 예측을 불러오는 것이 됩니다. 어떤 이벤트도, 자기 상관 회귀 모델, 또는 ARIMA 모델은 보다 좋게 적용 (리스크 시뮬레이터 (Risk Simulator) | 예측(Forecasting) | ARIMA)되지 않으면 않습니다. 최후에, 시리즈의 자기 상관을 관측하는 비계절성으로, 매끄럽게 하강하는 경향을 가지고 있습니다 (모델에서 비계절성 레포트를 참조하여 주십시오).

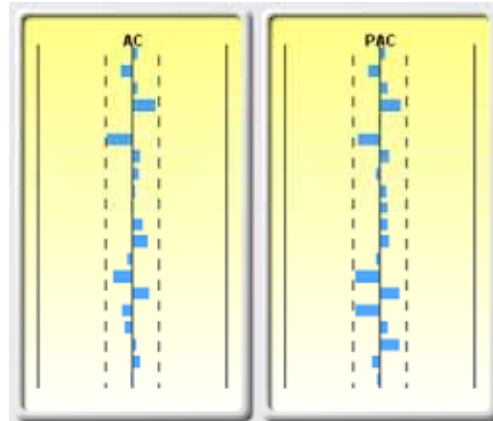
만약, 자기 상관 $AC(k)$ 이 제로와 같지 않은 경우, 이것은 시리즈가 먼저 최초에 연속적으로 상관되는 것을 표시하고 있습니다. 만약, $AC(k)$ 가, 증가하는 시간차와 기하학적으로 조금 떨어져 있는 경우, 시리즈는 저순위의 자기 회귀 과정을 따라 가는 것을 표시합니다. 만약, $AC(k)$ 가, 적은 수의 시간차의 후에 제로로 하강한 경우, 시리즈는 저순위의 이동 평균 과정을 따라 가는 것을 표시합니다. 일부의 상관 $PAC(k)$ 은, 개입한 시간차에서 상관을 소거한 후의 k 주기의 값의 상관을 측정합니다. 자기 상관의 패턴은, k 보다 적은 순위의 자기 회귀에 의하여 얻는 것이 가능, 시간차 k 에의 일부의 자기 상관은, 제로에 가까울 것입니다. Ljung-Box Q-통계와 시간차 k 에 대하여 이것의 p -값은, 귀무가설을 가지고, k 순서를 상회하는 자기 상관은 없습니다. 자기 상관의 플롯의 점선은, 근이적인 두 개의 표준 에러의 범위입니다. 만약 자기 상관이 이것의 범위를 가지고 있지 않은 경우, 제로에서 유의한 차이가 없고, 5%의 유의성 레벨을 표시하고 있습니다.

자기 상관은, 종속하고 있는 Y 변수 자신의 과거에의 관계를 측정합니다. 분포적인 시간차는 역으로, 종속하고 있는 Y 변수와 다른 독립하고 있는 X 변수의 사이의 시간차의 관계입니다. 예를 들어, 사망률의 분산과 방향은, 시간차(일반적으로 1에서 3개월)에 대하여 연방 펀드의 비율을 따라가는 경향을 보여 줍니다. 때에 따라 시간차는 싸이클과 계절성 (예, 아이스 크림의 판매는 여름에 증가하는 경향을 보임으로, 과거 12개월의 최후의 여름의 판매에 관련됩니다)을 따라 갑니다. 분산된 시간차 분석은(Figure 5.24), 여러 가지 시간차에 대하여 어떻게 종속 변수가 각 독립 변수에 관련 되어 있는가, 어떻게 시간차가 통계적으로 유의하게 고려되어 있는가를 정하기 위하여, 언제 모든 시간차로 정해져 있는가를 표시합니다.

가 2 ,2 1 ...) 가 (3
가 .() ARIMA .(I I ARIMA).
AC(1) , AC(k)가 , AC(k)
PAC(k) , k , k
가 가 , Ljung-Box Q- p k
5%
Y 가 (1 3)가 , X
12 . 가

Autocorrelation

Time Lag	AC	PAC	Lower Bound	Upper Bound	Q-Stat	Prob
1	0.0580	0.0580	-0.2828	0.2828	0.1786	0.6726
2	-0.1213	-0.1251	-0.2828	0.2828	0.9754	0.6140
3	0.0590	0.0756	-0.2828	0.2828	1.1679	0.7607
4	0.2423	0.2232	-0.2828	0.2828	4.4865	0.3442
5	0.0067	-0.0078	-0.2828	0.2828	4.4890	0.4814
6	-0.2654	-0.2345	-0.2828	0.2828	8.6516	0.1941
7	0.0814	0.0939	-0.2828	0.2828	9.0524	0.2489
8	0.0634	-0.0442	-0.2828	0.2828	9.3012	0.3175
9	0.0204	0.0673	-0.2828	0.2828	9.3276	0.4076
10	-0.0190	0.0865	-0.2828	0.2828	9.3512	0.4991
11	0.1035	0.0790	-0.2828	0.2828	10.0648	0.5246
12	0.1658	0.0978	-0.2828	0.2828	11.9466	0.4500
13	-0.0524	-0.0430	-0.2828	0.2828	12.1394	0.5162
14	-0.2050	-0.2523	-0.2828	0.2828	15.1738	0.3664
15	0.1782	0.2089	-0.2828	0.2828	17.5315	0.2881
16	-0.1022	-0.2591	-0.2828	0.2828	18.3296	0.3050
17	-0.0861	0.0808	-0.2828	0.2828	18.9141	0.3335
18	0.0418	0.1987	-0.2828	0.2828	19.0559	0.3884
19	0.0869	-0.0821	-0.2828	0.2828	19.6894	0.4135
20	-0.0091	-0.0269	-0.2828	0.2828	19.6966	0.4770



Distributive Lags

P-Values of Distributive Lag Periods of Each Independent Variable

Variable	1	2	3	4	5	6	7	8	9	10	11	12
X1	0.8467	0.2045	0.3336	0.9105	0.9757	0.1020	0.9205	0.1267	0.5431	0.9110	0.7495	0.4016
X2	0.6077	0.9900	0.8422	0.2851	0.0638	0.0032	0.8007	0.1551	0.4823	0.1126	0.0519	0.4383
X3	0.7394	0.2396	0.2741	0.8372	0.9808	0.0464	0.8355	0.0545	0.6828	0.7354	0.5093	0.3500
X4	0.0061	0.6739	0.7932	0.7719	0.6748	0.8627	0.5586	0.9046	0.5726	0.6304	0.4812	0.5707
X5	0.1591	0.2032	0.4123	0.5599	0.6416	0.3447	0.9190	0.9740	0.5185	0.2856	0.1489	0.7794

Figure 5.24 – 자기 상관과 분포적 시간차의 결과

회귀모델을 실행하는 것에 있어서 다른 요구는, 정규성의 가정과 에러의 구형입니다. 정규성의 가정이 위반, 및 이상치가 존재하는 경우, 선형 회귀의 적합의 이점은, 별로 강력하지 않고, 또한 가능한 정보 검정이 아니게 되어 버렸습니다. 그래서, 이것은 선형의 적합, 또는 그렇지 않은 것의 상이의 상위를 의미합니다. 에러가 독립하지 않고, 정규 분포되어 있지 않은 경우, 데이터는 자기 상관되지 않으면 안되고, 또는 비선형에서, 다른 무엇보다도 파괴적인 에러에서 영향되지 않은 것을 표시하고 있습니다. 에러의 독립성은 불균일 분산의 검정으로 검출할 수 있습니다(Figure 5.25).

실행된 에러의 정규성 검정은, 논파라메트릭 검정에서, 샘플이 그려진 모집단의 특정한 형상에 대하여 가정을 작성하지 않고, 작은 샘플 데이터의 설정의 분석을 인정합니다. 이 테스트는, 정규 분포된 모집단에서 그려진 샘플 에러가 어느 것인가 귀무 검정을 평가하는 것에 대하여, 대립 가설의 데이터 샘플은 정규 분포되어 있지 않습니다. 만약, 계산된 D 통계가, 여러 가지 유의한 값에 대하여 D-크리티컬의 값과 비슷하나, 상회하고 있는 경우, 귀무 가설을 방치하고, 대립 가설 (에러는 정규 분포되어 있지 않습니다) 를 인정합니다. 또한, D-통계가 D-크리티컬한 값보다 작은 경우는, 귀무 가설을 방치(에러는 정규 분포되고 있습니다)하지 않습니다. 이 검정은 두 개의 축적된 주파수에 영향을 불러 옵니다: 하나는 샘플 데이터 설정에서 얻고, 또 하나는, 샘플의 평균값과 표준 편차에 기초한 이론적 분포에서 얻을 수 있습니다.

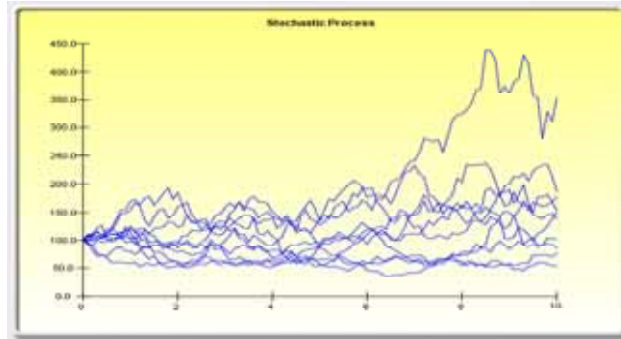
		O-E					
	0.00						
	141.83	-219.04	0.02	0.02	0.0612	-0.0412	
D	0.1036	-202.53	0.02	0.04	0.0766	-0.0366	
1% 시 D 임계치	0.1138	-186.04	0.02	0.06	0.0948	-0.0348	
5% D	0.1225	-174.17	0.02	0.08	0.1097	-0.0297	
10% 시 D 임계치	0.1458	-162.13	0.02	0.10	0.1265	-0.0265	
가 :		-161.62	0.02	0.12	0.1272	-0.0072	
		-160.39	0.02	0.14	0.1291	0.0109	
		-145.40	0.02	0.16	0.1526	0.0074	
: 1%		-138.92	0.02	0.18	0.1637	0.0163	

Figure 5.25 – 에러의 정규성 검정

때때로, 어떤 타입의 시계열 데이터는, 확률 과정 이외의 방법으로 모델화 할 수 없습니다. 이것은, 선으로 끌어낸 이벤트의 성질은 확률적이기 때문입니다. 예를 들어, 모델에 적합하지 않은, 주식 가격, 이윤, 석유의 가격과 물가를 심플한 회귀 모델을 사용하여 예측할 수 있다고 합시다. 이것은, 이러한 변수는 무엇보다도 불확실하고, 휘발하고있어, 먼저 정해진 동작의 정적인 룰을 따라 가지 않기 위함입니다. 즉, 과정은 계절성이 아닌 것입니다. 계절성은 검정의 실행의 사용에 의하여 확인되는 한편, 다른 시각계구는 자기 상관의 레포트로 검출됩니다 (ACF 는, 천천히 하강의 경향을 보여 줍니다). 확률 과정은, 이벤트의 연속, 및, 확률 법률에 의하여 생성되는 경로입니다. 이 랜덤 이벤트는, 시간이 지나 가면 발생할 수 있으나, 특정의 통계적, 그리고 확률적 룰로 지배됩니다. 주요 확률 결과는, 랜덤 워크, 혹은 브라운 운동, 평균 회귀와 샤프 확산이 포함되어 있습니다. 이것의 과정은 랜덤 경향을 따라서, 확률 법칙으로 정해진 복수의 변수를 예측하는 것에 사용됩니다. 과정을 생성하는 방식은 잘 알려져 있습니다만, 현재 생성된 결과는 알려져 있지 않습니다. (Figure 5.26).

랜덤 워크 브라운 운동의 과정은, 주식 가격, 물가와 다른 확률적 시계열 데이터가 부여하는 상향 조정, 또는 성장률과 상향 조정 경로 주변의 성장률을 예측하는 것에 사용할 수 있습니다. 평균 회귀의 과정은, 장기값을 목적으로 한 경로를 인정하고, 이윤 및 인플레이션의 비율 (시장 단속 권한에 의한 장기 목적) 과 같은 장기

비율을 가진 시계열의 변수의 예측을 위하여 사용하기 쉽고, 장기 목적의 비율의 랜덤 워크 가정의 분산을 감소할 수 있습니다. 점프 J 가정은, 석유의 가격, 또는, 전기의 가격 (분리된 외인성의 이벤트의 충돌은 가격을 상승하거나 하강합니다) 과 같은 변수가 랜덤 점프를 때대로 표시할 수 있을 때에 시계열 데이터의 예측에 적용합니다. 이것의 가정은 필요에 따라서 혼합하여 적합한 것이 가능합니다.



가		, Jump-diffusion		()	
가		()		가	
-1.48%		283.89%		Jump 20.41%	
88.84%		327.72		Jump 237.89	
		46.48%			
Runs	20		-1.7321		
	25	P-Value (1-tail)	0.0416		
	25	P-Value (2-tail)	0.0833		
Run	26				
p	(0.10, 0.05, 0.01)				, ARIMA
	p				

Figure 5.26 - 확률 과정 파라미터의 추정

다중 공선성의 존재는, 선형적 관계가 독립 변수 사이에 있을 때에 보여 집니다. 이것이 일어날 때는, 회귀 방식은 완전히 추정할 수 없습니다. 공선성에 가까운 상황에서는, 추정된 회귀의 방식은 바이어스 되어, 불확실한 결과를 부여합니다. 이 상황은, 스텝 넓이의 회귀가 적용되는 때에 특히 발생하고, 통계적으로 유의한 독립 변수는, 예상보다 빨리 회귀 믹스에 던져 버리고, 회귀 방식의 결과로서, 어느 쪽도 효과적으로 확실합니다. 중회귀 방식의 다중 공선성의 존재의 하나로써 바른 검정은, R-평방 된 값은 비교적 높고, t-통계는 비교적 낮다는 것입니다.

더 하나의 신속한 검정은, 독립 변수 사이의 상관 매트릭스를 작성하는 것입니다. 높은 횡단적 상관은, 자기 상관의 강함을 표시합니다. 눈대중으로는, 상관과 0.75 보다 큰 절대치로, 엄격한 다중 공선성의 표입니다. 다중 공선성을 위하여 다른 검정은, 분산

인플레이션 요소 (VIF)의 사용으로, 각 독립 변수를 모든 다른 독립 변수에 회귀하여 얻는 것입니다. R- 평방의 값과 VIF 의 | 계산을 얻는 것이 가능합니다. 2.0 를 넘는 VIF 는, 엄격한 다중 공선성으로서 고려되고 있습니다. 10.0 을 넘는 VIF 는, 파국적인 다중 공선성을 표시하고 있습니다 (Figure 5.27)。

상관관계 매트릭스					
CORRELATION	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	
X5				1.000	

편차 Inflation 요소					
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	
X5				N/A	

Figure 5.27 - 다중 공선성의 예러

상관 매트릭스는, 변수의 페어 사이의 피어슨 상공 모멘트 상관 일반적으로 피어슨의 R 로서 참조되어 있음)으로서 표시되고 있습니다. 상관 계수는, -1.0 과 + 1.0 사이의 범위에 포함되어 있습니다. 부호는, 변수 사이의 관계의 방향을 표시하고, 계수는 관계의 강함과 크기를 표시합니다. 피어슨 R 은, 선형적 관계를 측정하고, 비선형의 관계의 측정에서는 보다 적은 효과를 표시합니다。

어느쪽의 상관이 유의한가를 검정하는데는, 양측 테일의 가설 검정이 실행되고, 결과가 된 p-값은 상기의 리스크에 표시되어 있습니다. 0.10, 0.05, 과 0.01 보다 작은 p-값은, 파란색으로 표시되어 있고, 통계적인 유의성을 표시하고 있습니다. 즉, 상관의 페어를 위하여 부여된 유의한 값보다 작은 p-값은, 제로에서 통계적으로 달라서, 두 개의 변수의 사이의 유의한 선형이 관계가 있는 것을 표시합니다。

두 개의 변수 (x 와 y)의 사이의 피어슨 상공의 일차 관계수(R)은, 공분산(cov)의 측정에 관련 되어, $R_{x,y} = \frac{COV_{x,y}}{s_x s_y}$ 로 표시되어 있습니다. 공분산을 두 개의 변수의 표준

편차(s)로 나누는 것의 이점은, 결과가 되는 상관계수는 -1.0 과 +1.0 의 사이에 해당됩니다. 이것은 상관을 좋은 비교 측정으로서, 다른 변수 사이의 비교 (특히 다른

Unit 과 크기)를 부여해 주고 있습니다. 스피어맨의 비선형 상관에 기초한 순위도 하기에 포함되어 있습니다. 스피어맨 R 은, 피어슨의 R 에 관련되어 있고, 데이터는 먼저 순위별로 랭크되어, 그 후 상관됩니다. 랭크의 상관은, 두 개의 변수의 하나, 또는 양방과 함께 비선형의 때에, 이것의 관계는 보다 좋게 추정됩니다.

유의한 상관은 관련 되지 않는 것으로 중점을 두어 주십시오. 변수간의 관계는, 한 개의 변수는 다른 변수의 변환을 불러일으키는 원인이 됩니다. 두 개의 변수가 서로 단독적으로 이동하고 있습니다만, 관련서의 경로에서는 이것은 상관되어 있으나, 이것의 상관 관계는 비슷하지 않으면 않습니다(예, Sun Spot 과 주식 마켓 사이의 상관은 강해지지 않으면 않되지만, 한 개는 추정 가능하고, 원인이 없고, 이것의 상호 관계는 순수하게 비슷합니다).

다른 강력한 리스크 시뮬레이터의 틀은 통계 분석 툴로서, 데이터의 통계적인 성질을 정의합니다. 진단의 실행은, 데이터의 확률 성질과 검정을 위하여 기본적인 통계의 기술에 의하여 복수의 통계적 성질의 데이터 확인이 포함되어 있습니다.

순서:

- 샘플 모델을 열고(리스크 시뮬레이터 (Risk Simulator) | 샘플 (Examples) | 통계적 분석 (Statistical Analysis)) , 데이터 워크 시트와 변수 명칭을 포함한 데이터를 선택하여 주십시오 (셀 C5:E55).
- 리스크 시뮬레이터 (Risk Simulator) | 툴 (Tools) | 통계분석 (Statistical Analysis) (Figure 5.28) 을 선택하여 주십시오.
- 데이터 타입을 확인하고, 선택된 데이터는, 열에 변환된 단일 변수에서이거나 아니면 복수의 변수에서 등 어느쪽인가 입니다. 이 예증에서는, 선택된 데이터의 범위는, 복수의 변수에서라고 고려하여 주십시오. 종료 후, OK 를 클릭하여 주십시오.
- 희망하는 통계적 검정을 선택하여 주십시오. 추천은 (디폴트로) , 모든 검정을 선택하는 것입니다. 종료 후, OK 를 클릭하여 주십시오(Figure 5.29).

실행된 검정의 보다 좋은 독해를 위하여 생성된 레포트를 보시기 바랍니다 (Figures 5.30-5.33 로 샘플 레포트가 표시되어 있습니다).

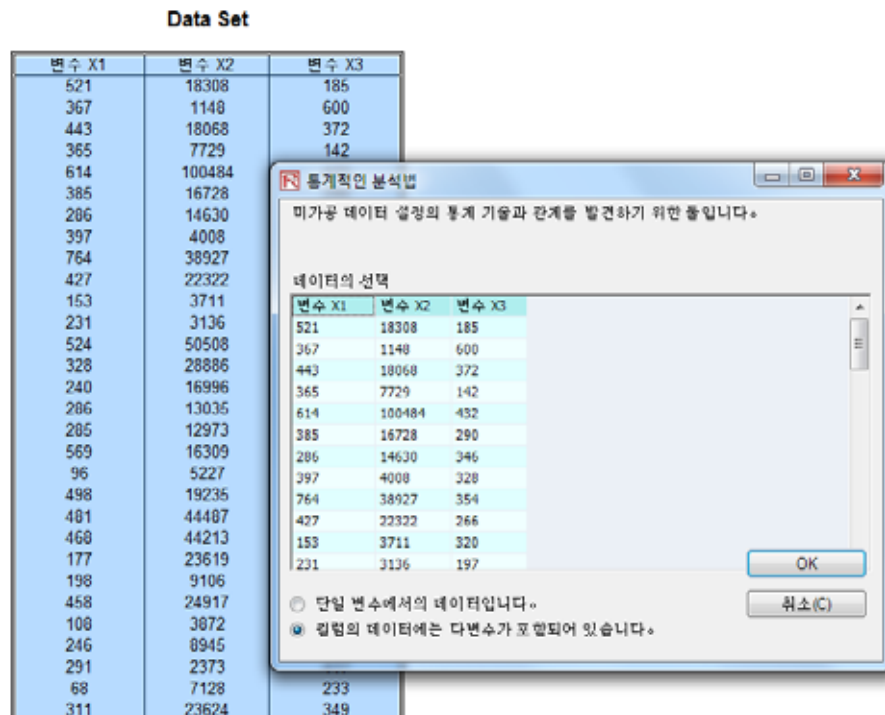


Figure 5.28 -통계적 분석 툴의 실행

Statistics from Dataset:		t-Statistic	
Sample	50	t-Statistic	13.5734
Sample	331.92	P-Value (right-tail)	0.0000
Sample	172.91	P-Value (left-tailed)	1.0000
		P-Value (two-tailed)	0.0000

가 (Ho): $\mu = 74$

가 (Ha): $\mu < 74$

Notes: "<"

가

t 가 가). t 가 가 (t : 가 ,

가 가 .

가

가 가 Ho 가 (0.10, 0.05, 0.01) ,

가 10%, 5% 1% (90%, 95%, 99%) , p 0.10, 0.05, 0.01 ,

가 ,

가

가 가 Ho 가 (0.10, 0.05, 0.01) ,

가 10%, 5% 1% (90%, 95%, 99%) , p 0.10, 0.05, 0.01 ,

가

가 가 Ho 가 (0.10, 0.05, 0.01) ,

가 10%, 5% 1% (90%, 95%, 99%) , p 0.10, 0.05, 0.01 ,

가

왜냐하면, t -검정은 보다 보수적이고 Z -검정에서처럼 알려진 모집단 표준편차가 필요치 않으므로 우리는 이 t -검정만 사용한다.

Figure 5.31 -통계적 분석 툴의 레포트 샘플(한 개의 변수의 가정 검정)

Figure 5.32 -통계적 분석 툴의 레포트 샘플(정규성 검정)



Figure 5.33 - 통계적 분석 툴의 샘플 레포트(확률 파라미터의 추정)

이것은, 리스크 시뮬레이터에서의 통계적 확률 툴로써, 다양한 설정으로 유용하고, 확률 밀도 관수 (PDF)의 계산이 이산 분포 (이것의 명칭을 혼합하여 사용합니다)에서는 확률 관수 (PMF)에도 사용할 수 있고, 몇 개의 분포와 파라미터가 부여되어 있기 때문에, 어떤 결과 x 의 발생 확률을 정하는 것이 가능합니다. 또한, 누계 분포 관수(CDF)는 계산할 수 있고, 이 x 값까지의 PDF 의 덧셈입니다. 최후에, 반누계분포 관수(ICDF)는, 소여의 발생으로 확률 x 값의 계산에 사용할 수 있습니다.

이 툴은, 리크스 시뮬레이터 (Risk Simulator) | 툴 (Tools) | 분포적인 분석 (Distributional Analysis) 을 선택하는 것으로 기동할 수 있습니다. 샘플과 같이, Figure 5.34 에서는, 2 항 분포의 계산(예, 동전을 던지는 등 2 개의 결과를 가진 분포, 결과가 표 혹은 테일의 어느쪽인가로, 이것에 규정된 확률이 있음)을 표시하고 있습니다. 예를 들어, 동전을 2 회 던져서, 결과의 표를 성과로 설정하여, 2 항 분포를 시행 회수 = 2 (동전을 2 회 던진다)로 하여, 확률을 = 0.50 (표가 나오면 성공의 확률)로 합니다. PDF 의 선택과 x 값의 범위의 설정을 0 에서 2 로 하여, 스텝 사이즈를 1 (이것은, x 값에 0, 1, 2 을 요구하고 있는 것을 의미합니다)로 하여, 분포의 이론적 4 개의 모멘트와 같은 결과가 된 확률은 표와 그래프에 부여됩니다. 던진 동전의 결과로서 표-가 된 경우, 뒷면-뒷면, 표-뒷면과 뒷면-표로, 표에 나오지 않는 확률은 25%이 되고, 표 1 회는, 50%와 표 2 회는 25%입니다.

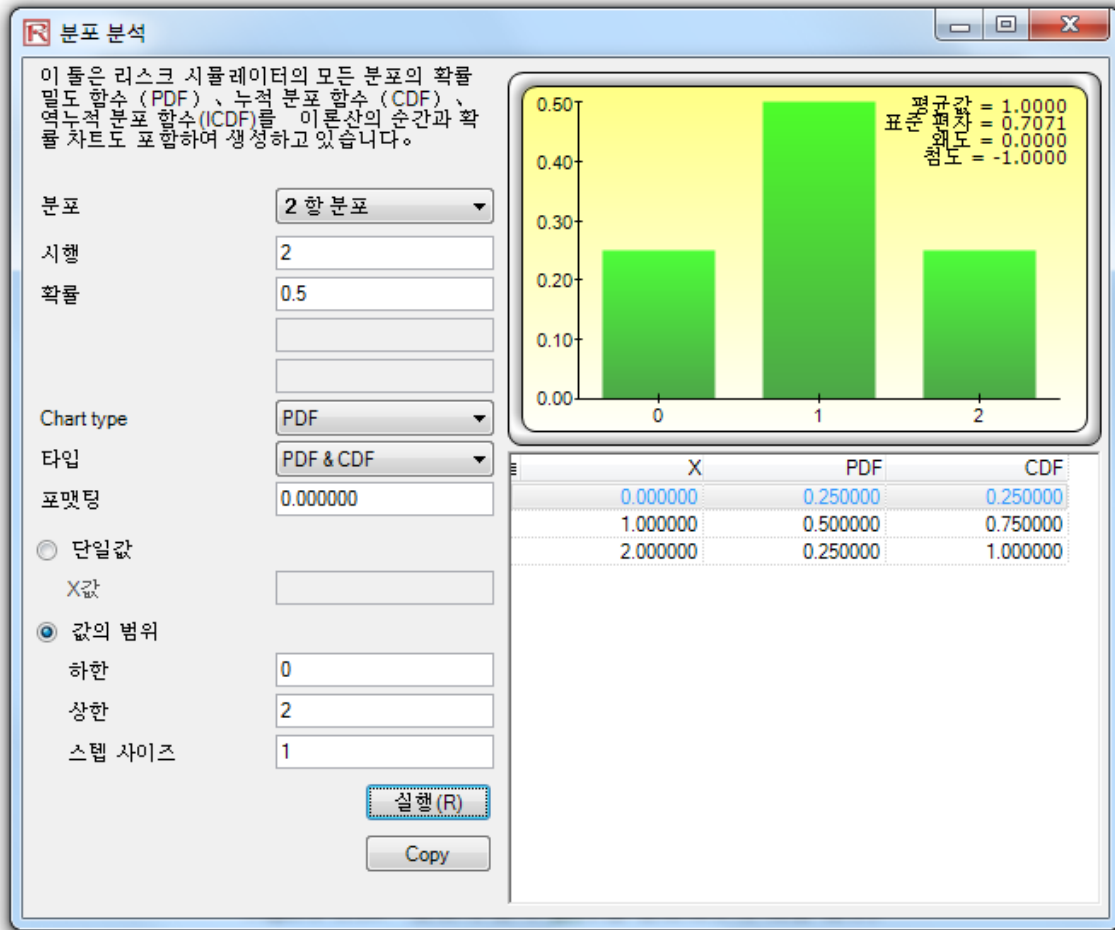


Figure 5.34 – 분포적 분석 툴(2 항 분포와 2 의 시행 회수)

같이, 동전을 던져서 엄격한 확률을 얻는 것이 가능합니다. Figure 5.35 로 표시되지 않은 것과 같이, 2 0 의 시행 회수를 행합니다. 결과는, 표과 그래프의 양방에 기술되어 있습니다.

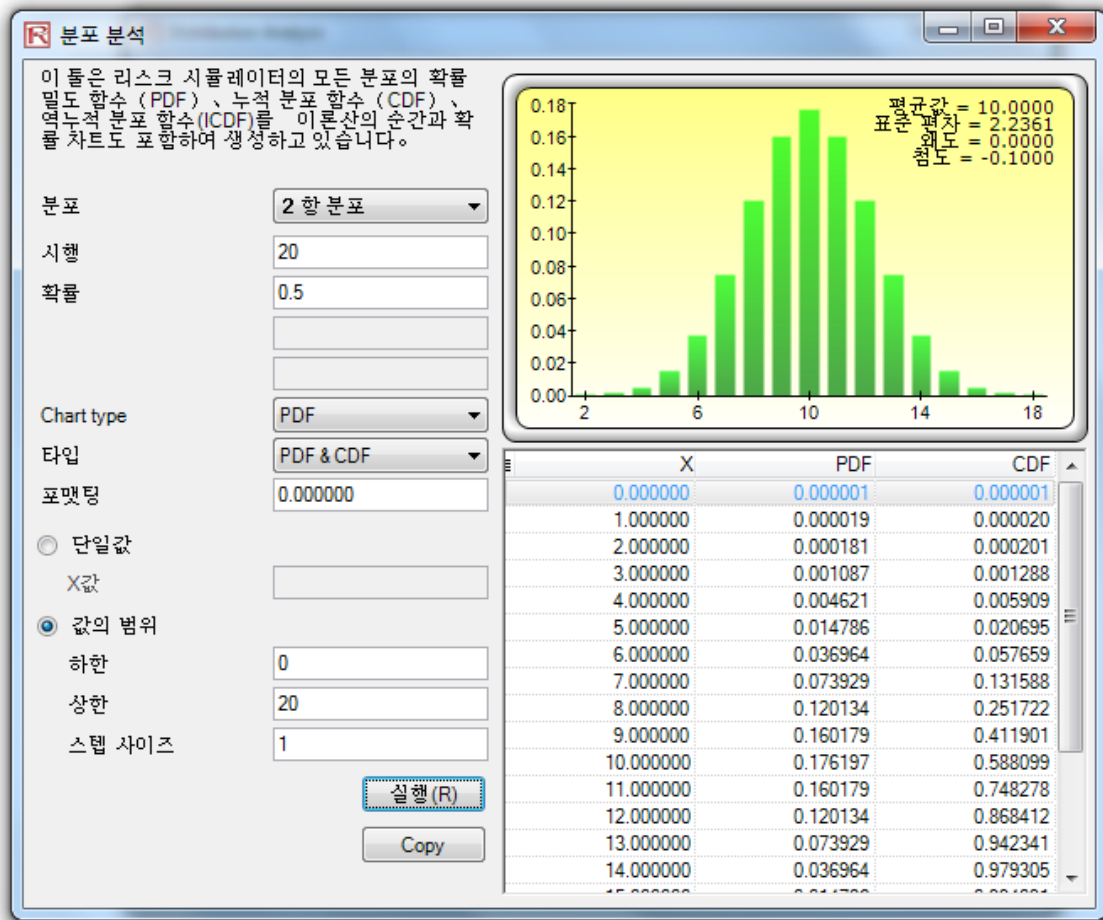


Figure 5.35 -분포적 분산 툴(2항 분포와 20의 시행 회수)

Figure 5.36 는, 같은 2항 분포를 표시하고 있으나, 여기에서는 CDF가 계산되어 있습니다. CDF란, 단순히 포인트 x 까지의 PDF의 값의 덧셈입니다. 예를 들어, Figure 5.35에서는, 0.1, 과 2의 확률은 0.000001, 0.000019, 과 0.000181이라는 것을 알 수 있고, 이것의 덧셈은 0.000201로, $x=2$ 의 CDF 값은 Figure 5.36로 표시되어 있습니다. PDF가 2개의 표를 얻는 것의 확률(및 0.1, 과 2의 표 확률)을 계산합니다. 분포는, (예, $1 - 0.00021$ 는 0.999799, 및 99.9799%를 얻음)은, 3개의 표, 및 그 이상을 얻는 확률을 부여하여 주고 있습니다.

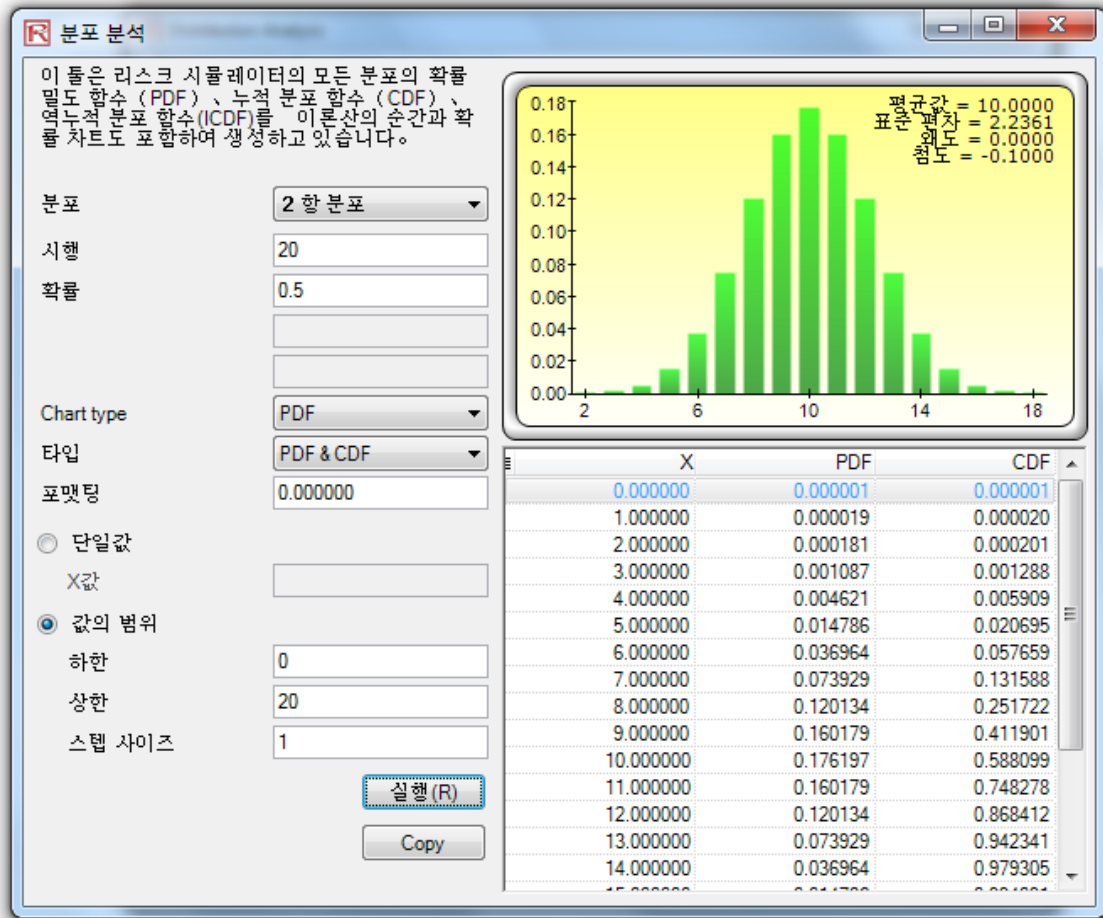


Figure 5.36 - 분포 분석 툴(2 항식 분포의 CDF 와 20 의 시행회수)

이 분포 분석 툴의 사용에 의하여, 리스크 시뮬레이터에 감마, 베타, 마이너스의 2 항분포와 같은 고도의 분포 분석이 가능합니다. 연속적인 분포와 ICDF 기능에서의 툴 사용의 샘플 이상으로, Figure 5.37 는 표준 정규 분포 (평균값이 제로로 표준편차 1 의 정규 분포) 를 표시하고, ICDF 의 적용에 의하여, 97.50% (CDF)의 누적 확률에 해당하는 x 의 값을 보여줍니다. 이것은, 97.50%의 절편 CDF 가, 양측 신뢰 구간의 95% (우측 확률로서 2.50%, 좌측의 확률로서 2.50%에 상당하고, 95%를 중심 및 신뢰 구간의 범위에 남아, 하나의 테일의 범위에 97.50%가 동등합니다) 에 비슷하다는 것입니다. 결과는, 잘 알려져 있는 1.96 의 Z-스코아입니다. 따라서, 분포적 분석 툴, 다른 분포를 위한 표준 스코아, 다른 분포의 누계확률, 및 일치의 사용은, 빠르고 또한 간단히 얻을 수 있습니다.

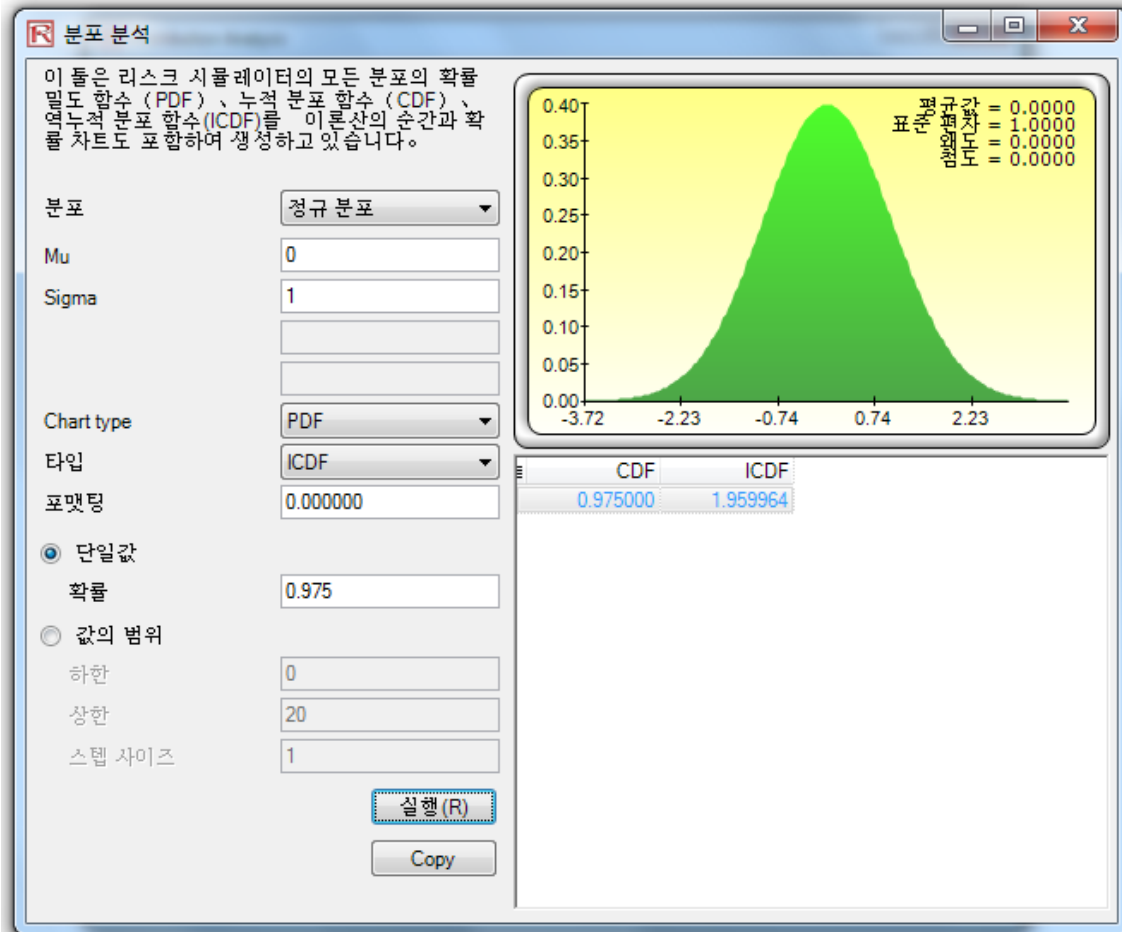


Figure 5.37 - 분포 분석 툴 (정규 분포 ICDF 와 Z-스코어)

Risk Simulator 2011/2012

,
 2001/2012 6 (Random Number Generators), 3 Correlation
 Copulas , 2 Simulation Sampling (5.41). **Risk**
Simulator | Options

Random Number Generator (RNG,)
 . ROV Risk
 Simulator proprietary 가
 6 가 가 ROV Risk Simulator default
 Advanced Subtractive Random Shuffle 2 가 가

.
 가
 , 가

Correlations Normal Copula, T-Copula, Quasi-Normal copula 3 가
 . ,
 Normal Copula 가 t-copula
 , quasi-normal copula

Monte Carlo Simulation (MCS) Latin Hypercube Sampling (LHS)
 . Copulas LHS
 LHS
 LHS , 가
 가 , 가
 , LHS . LHS

가
가
LHS
가
가
LHS
()
(가)

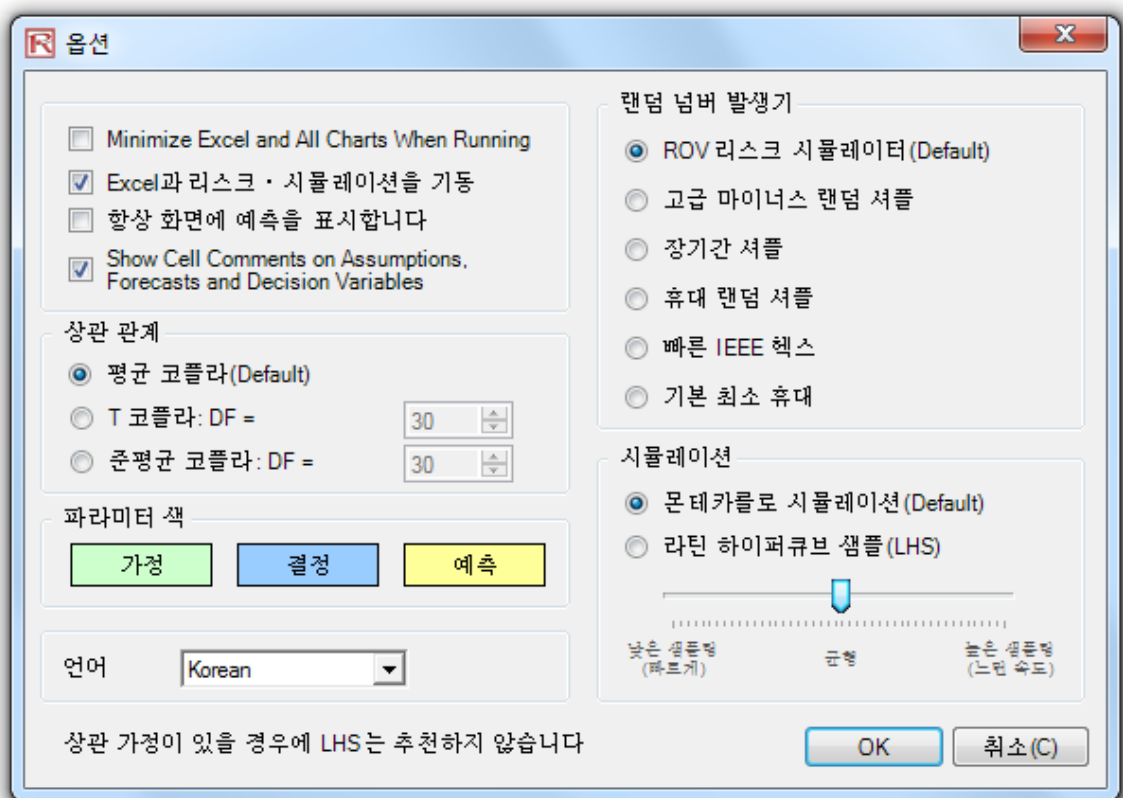


Figure 5.41 – Risk Simulator Options

Deseasonalize Detrend

deseasonalize

detrend (Figure 5.42).

가

(, ,).

(, 24 , 12 , 4 , 60 ,).

deseasonalize detrend . 가 1

1 .

(Deseasonalization and Detrending):

 (: B9:B28) ***Risk Simulator | Tools | Data***

Deseasonalization and Detrending

 ***Deseasonalize Data / Detrend Data*** ,

detrending , (: polynomial order, moving average order, difference order, and rate order) ***OK***

 가 , , , deseasonalized/detrended

(Deseasonalization and Detrending):

 (: B9:B28) ***Risk Simulator | Tools | Data***

Seasonality Test

 , 6

: 1, 2, 3, 4, 5, 6. 1

 , , ,

가 (가 RMSE
) : root mean
 squared error (RMSE), mean squared error (MSE), mean absolute deviation
 (MAD), mean absolute percentage error (MAPE)

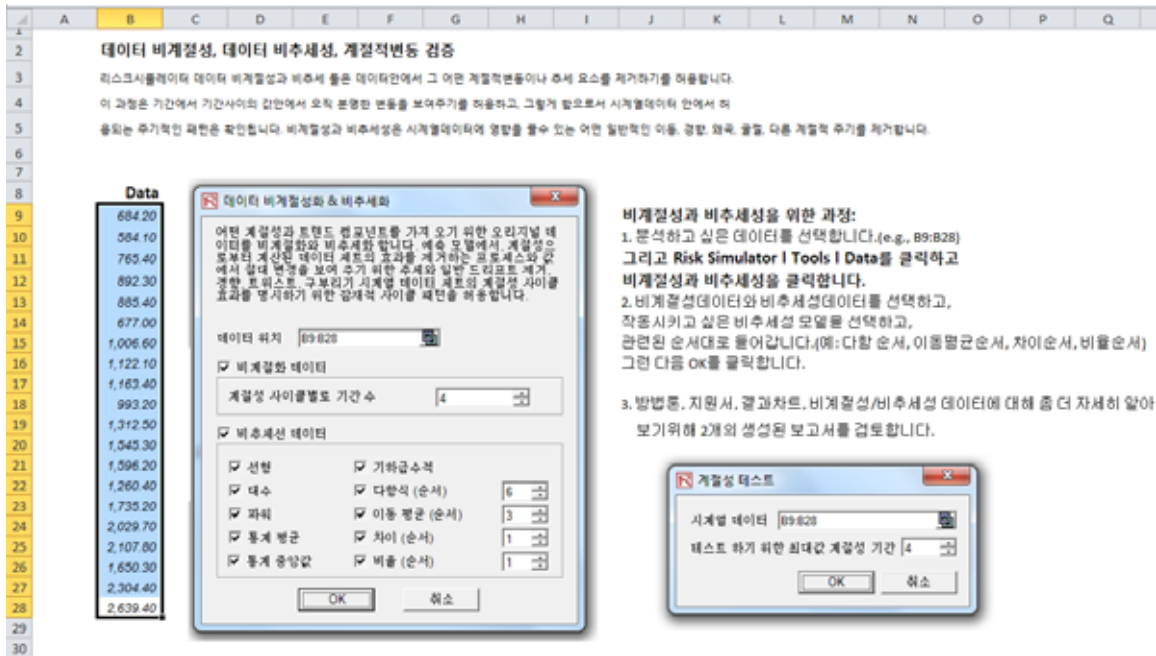


Figure 5.42 – Deseasonalization and Detrending Data

Principal Component Analysis

(Figure 5.43).

가 가 ,
 ,
 가 가

:



(: B11:K30), **Risk Simulator | Tools | Principal**

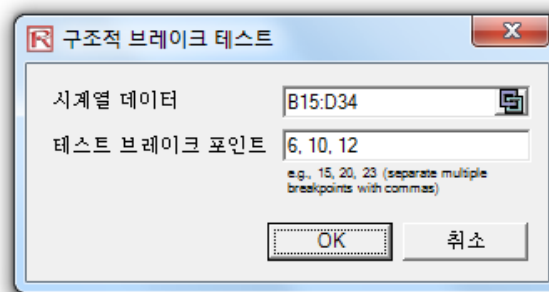
Component Analysis

OK



:
 (: B15:D34), **Risk Simulator | Tools | Structural**
Break Test
 , OK
 가
 ..

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



구조적 브레이크 테스트

시계열 데이터: B15:D34

테스트 브레이크 포인트: 6, 10, 12
e.g., 15, 20, 23 (separate multiple breakpoints with commas)

OK 취소

Figure 5.44 – Structural Break Analysis

Trendline

Trendline 가

(Figure 5.45).

(exponential,

logarithmic, moving average, power, polynomial, or power).

:



, **Risk Simulator | Forecasting | Trendline**

trendline (:

), (: 6 periods), OK

trendline 가 가

Year	Quarter	Period	Sales
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

트렌드 라인

추세선

☒ 선형
 ☒ 기하급수적
 ☒ 다항식 (순서) 2
 ☒ 이동 평균 (순서) 2

예측 적합 3 기간

OK

취소(C)

Figure 5.45 – Trendline Forecasts

Model Checking Tool

(Figure 5.46).

가 , Risk Simulator | Tools
 | Check Model . 가

Risk Simulator 가



Percentile Distributional Fitting tool (Figure 5.47)

가

- Distributional Fitting ()— (/)

가

가

가

가

- Distributional Fitting ()—가

42

가

- Distributional Fitting ()—

가

- Custom Distribution (가)—

:

 **Risk Simulator | Tools | Distributional Fitting (Percentiles)**

Run

R-square

가

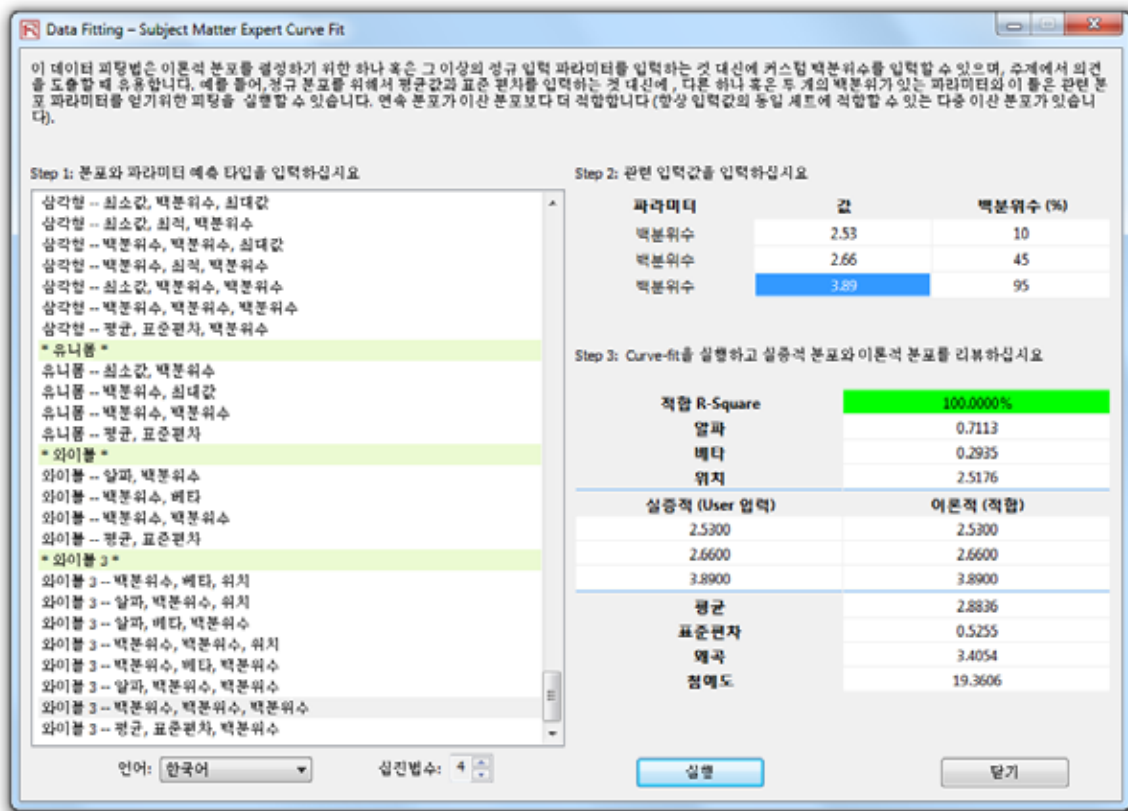


Figure 5.47 – Percentile Distributional Fitting Tool

:

(Figures 5.48-5.51). Risk Simulator 가

- Distributional Analysis— Risk Simulator 가 42 PDF, CDF, ICDF ,
- Distributional Charts and Tables— ,
(: [2, 2], [3, 5], [3.5, 8] 가 PDF, CDF, ICDF ,
).
- Overlay Charts— (가
(: [2, 2], [3, 5], [3.5, 8] 가
PDF, CDF, ICDF ,
).

Inputs

Run

CDF,

ICDF, PDF 가 45

(Figure 5.48).

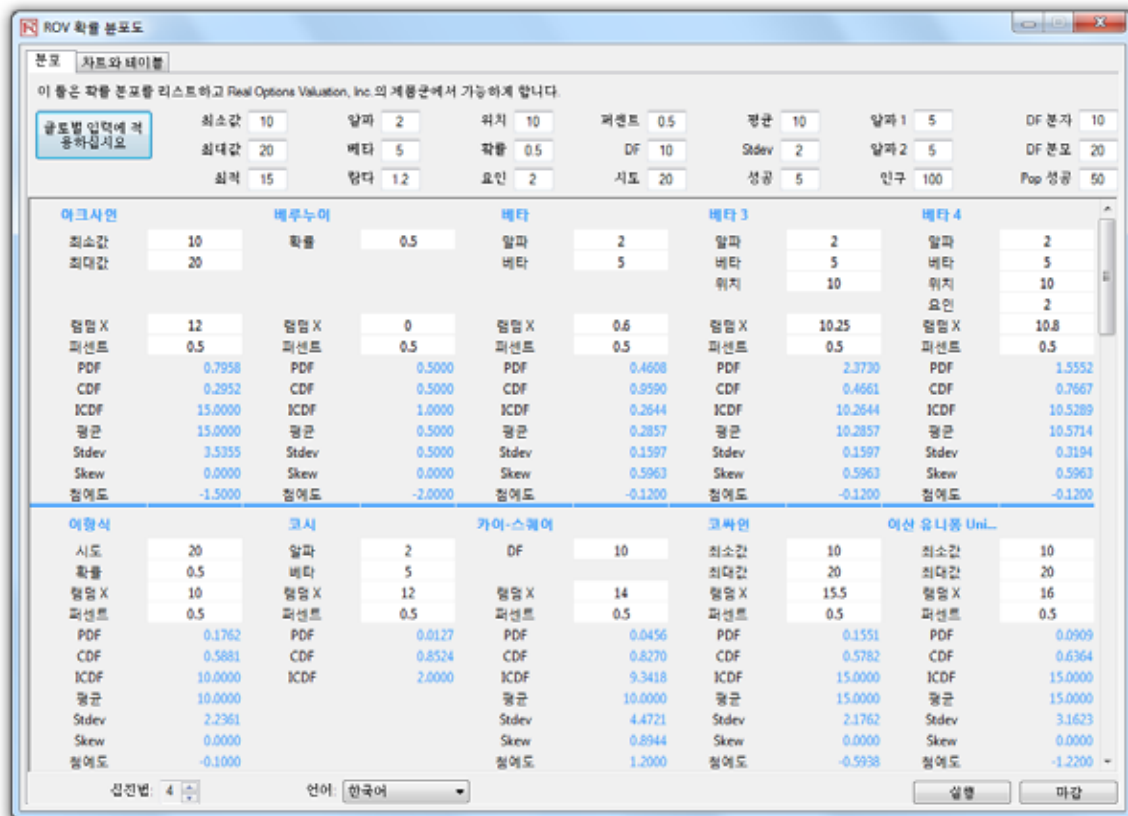


Figure 5.48 – Probability Distribution Tool (45 Probability Distributions)

Charts and Tables tab

(Figure 5.49),

[A] (:

Arcsine), CDF, ICDF, PDF

[B],

click **Run Chart** **Run Table**

[C].

[E]



From/To/Step

Custom

Run

[H]

Figure 5.50

PDF [G], [I]

[H], :

2;5;5, 5;3;5 [J] *Run Chart*

가 [K]: (2,5), (5,3), (5,5) [L].

[M],

[N].

Figure 5.51 (X)가 *From/To/Step*

[O],

Table tab [P]

0., 0.25, ..., 0.50 가 row variable

0, 1, 2, ..., 8 가 column

variable 가 = 20, = 0.5,

X = 10 , PDF 가

가 , 25%

가 20

0.0669 6.69% .

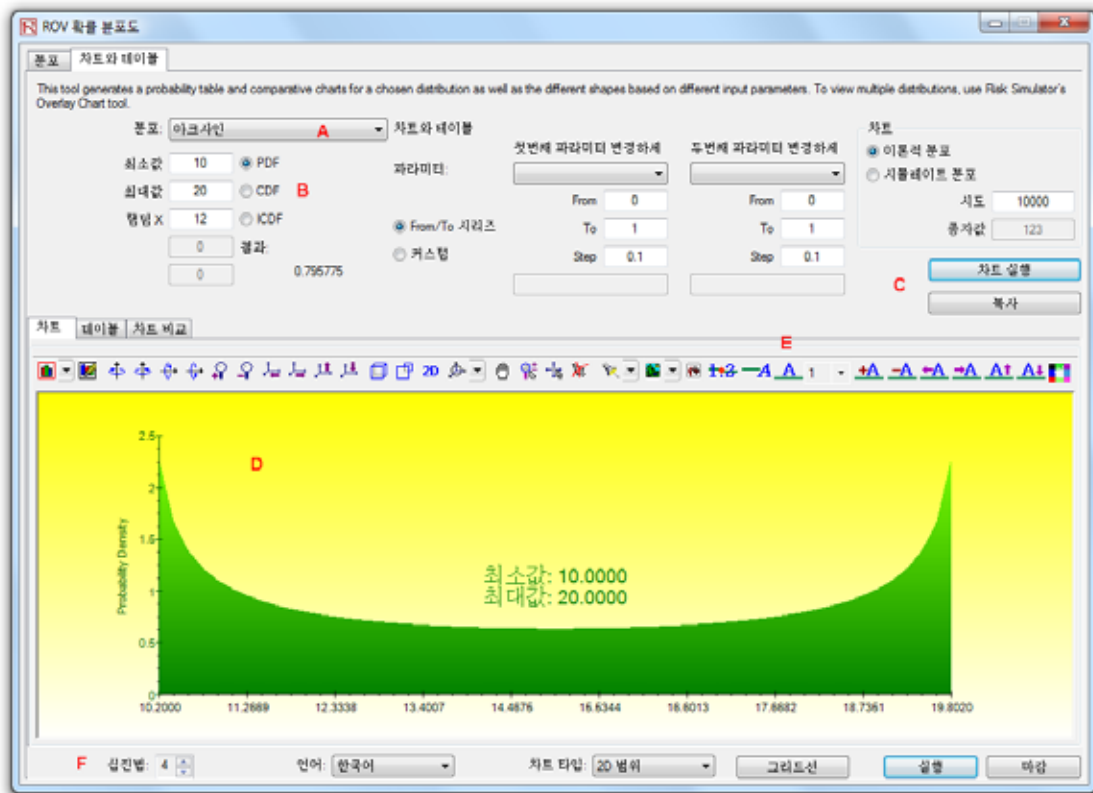


Figure 5.49 – ROV Probability Distribution (PDF and CDF Charts)

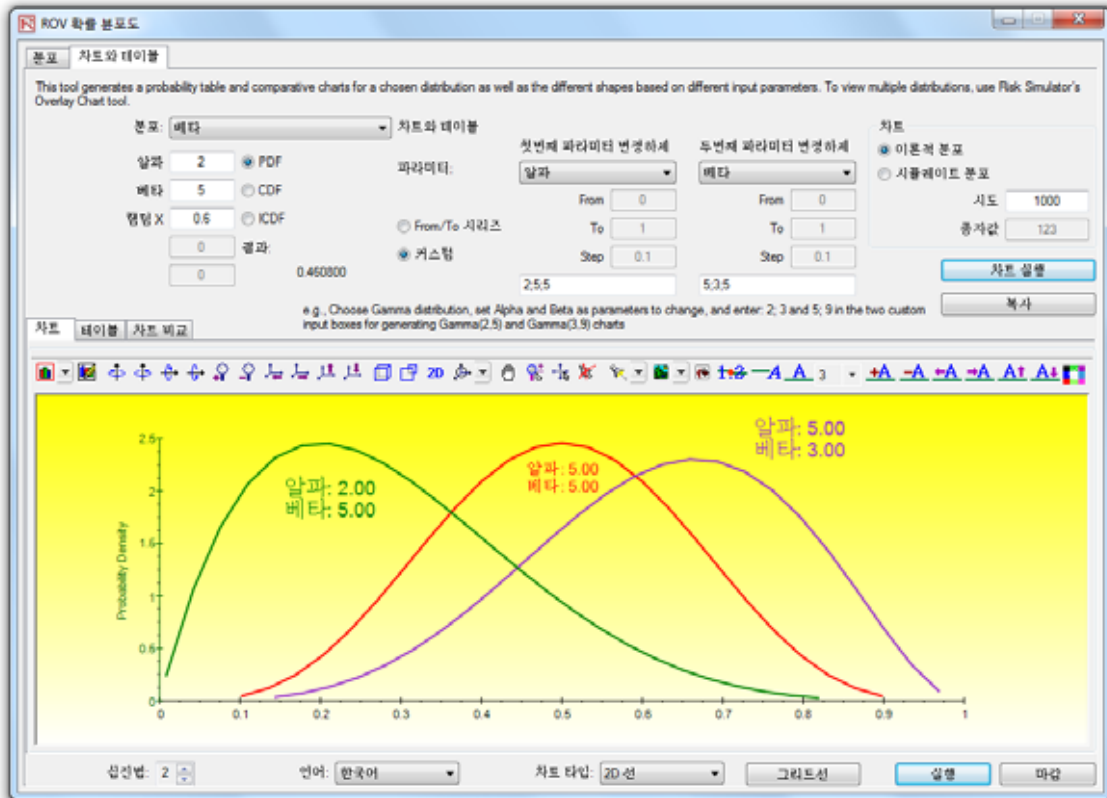


Figure 5.50 – ROV Probability Distribution (Multiple Overlay Charts)

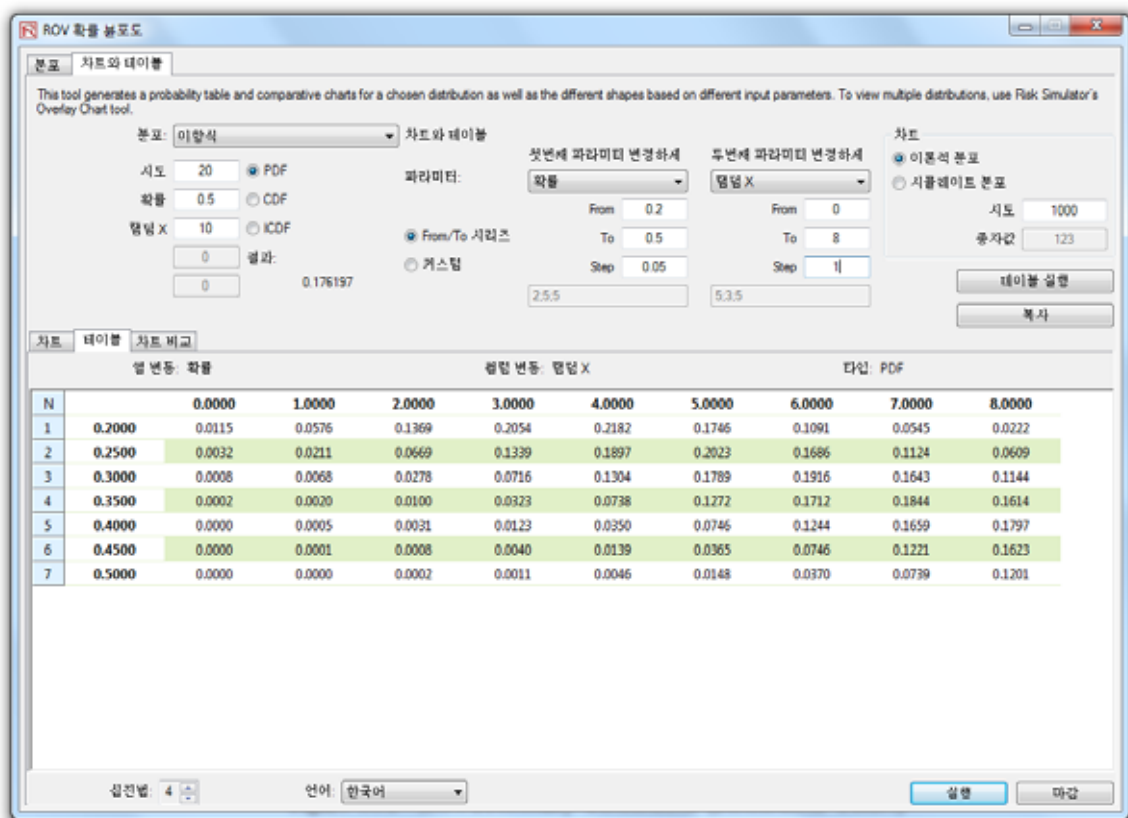


Figure 5.51 – ROV Probability Distribution (Distribution Tables)

ROV BizStats

ROV BizStats

Risk Simulator

, 130 가

(Figures 5.52-5.55).

:

 **Risk Simulator | ROV BizSta**, ROV BizStats ,

[A]

Example

/

[D] (Figure

5.52).

Notes row [C]

 Step 2

[F]

[G]

[H].

()



[J]

Run [I]

. Step 3

tab



Step 4 [L]

가 가 [M],

[N],

가 *.bizstats

[O]

:



1,000



[S]

[A]

가

[J]

[T/U]

[G]

- 가 [H] ,
- 가, , , , 가 [D] ,
- 가 [P]
- Command console [V/W/X].
- [S] Command console
- [V] , 가 Run
- Command [X] . Console
- *.bizstats ()
- XML XML [Z] 가 .
- 가
- Auto Fit, Cut, Copy, Delete, Paste
- Visualiz
- 가 (
- auto fit).
- Home End
- : Ctrl+Home (), Ctrl+End (),
- Shift+Up/Down ()
- Notes

- Visualize (: , , / , , 가)
- Copy 3 Results, Charts, Statistics copy
- Report 4 가
- Microsoft Excel , Microsoft PowerPoint 가
- Example 1 가 2 가 가
- Language 10 가 가 가
- 2 View drop list (: , ,) 가
- 가 가 (: 50 1,000.50 1,000,50 1'000,50). ROV BizStats [Data | Decimal Settings](#) English USA , ROV BizStats North America 1,000.50 (ROV BizStats).

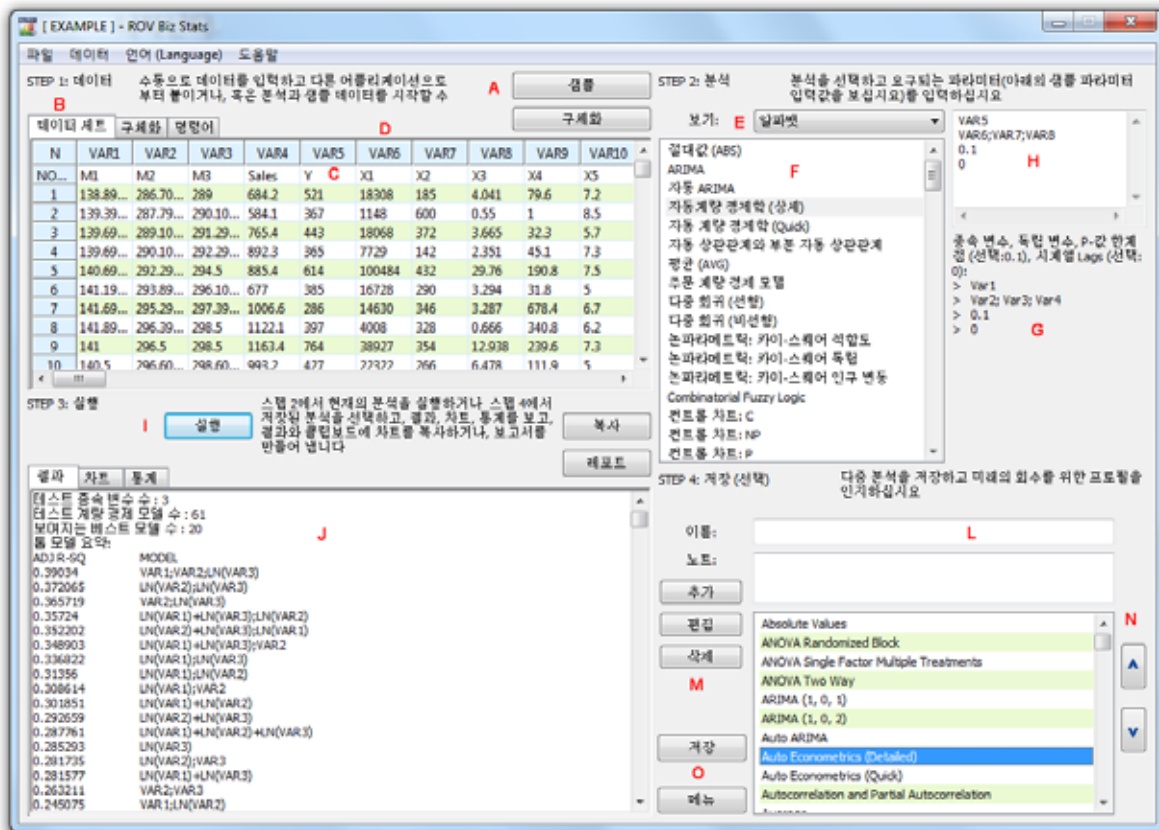


Figure 5.52 – ROV BizStats (Statistical Analysis)

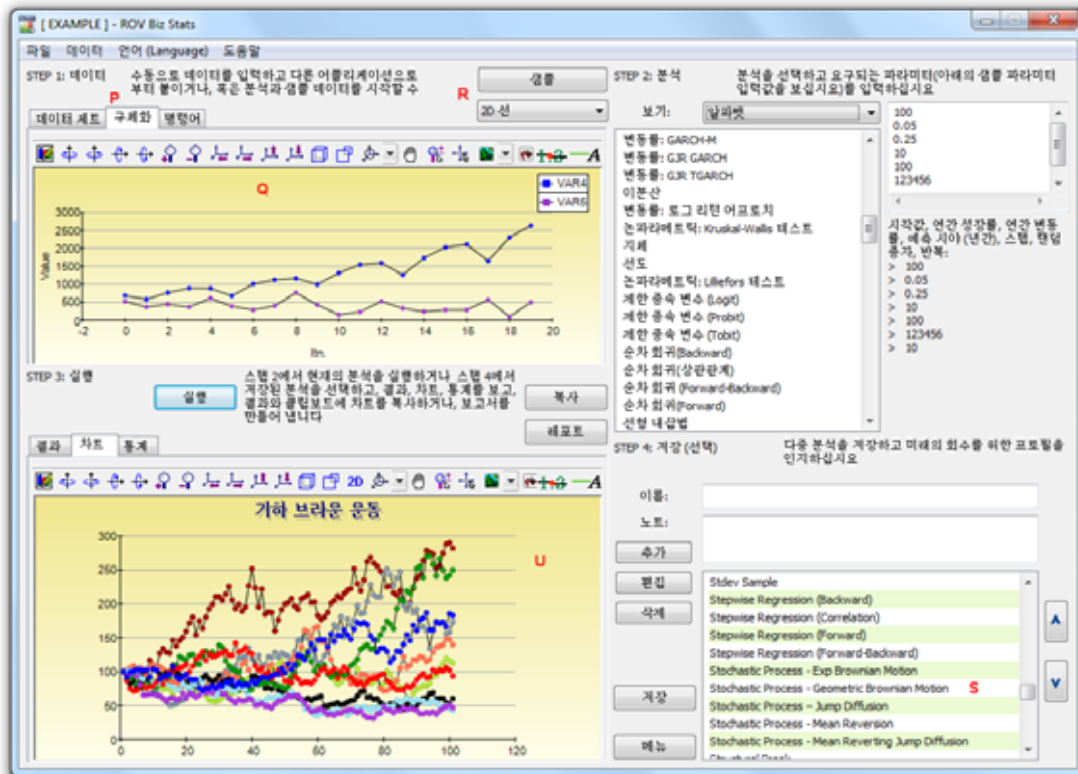


Figure 5.53 – ROV BizStats (Data Visualization and Results Charts)

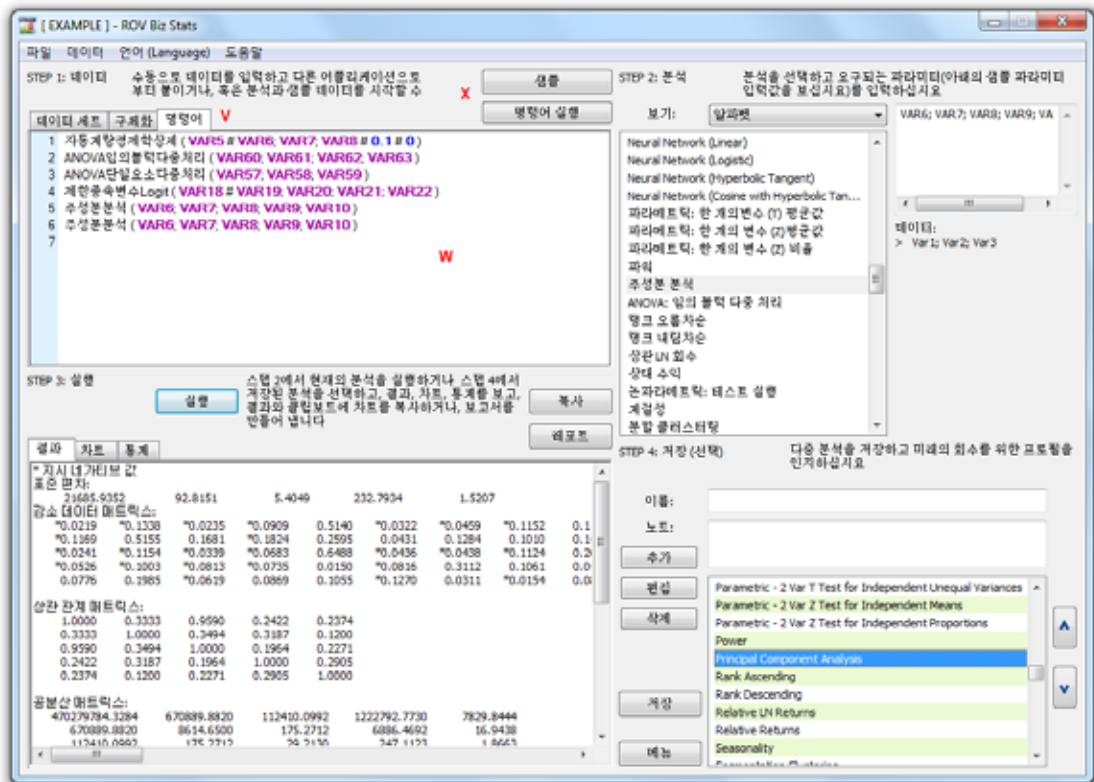


Figure 5.54 – ROV BizStats (Command Console)

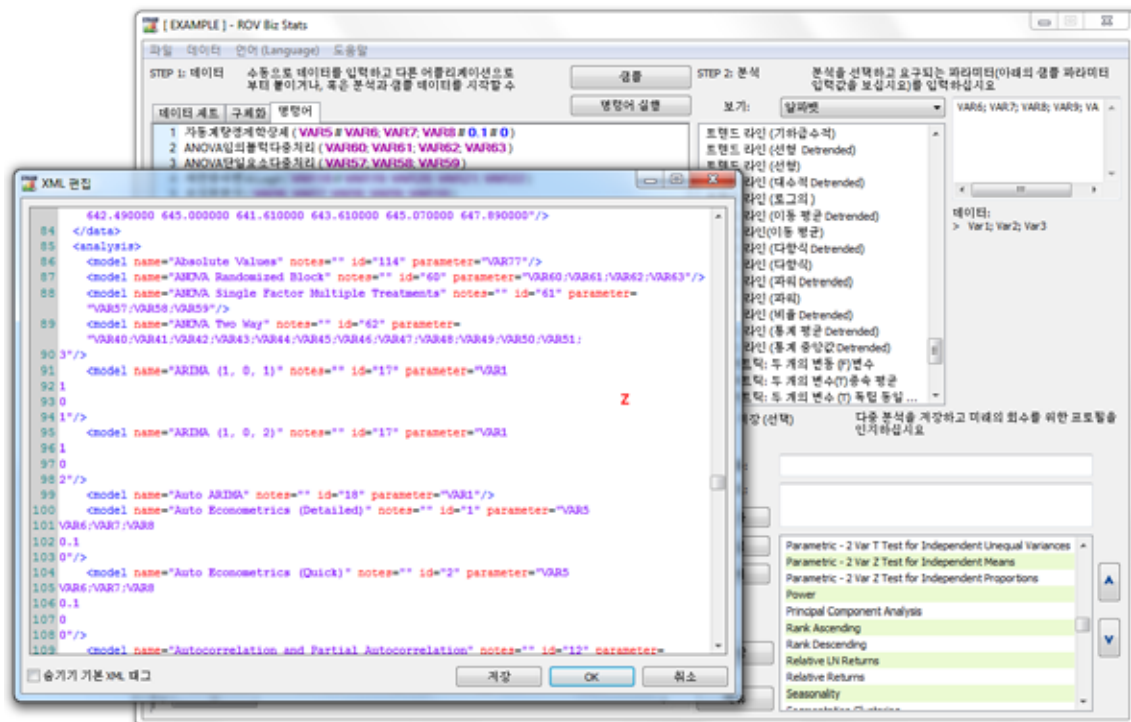


Figure 5.55 – ROV BizStats (XML Editor)

nodes)

Risk Simulator | ROV BizStats | Risk
 Simulator ROV BizStats 5.56

Risk Simulator |
 (가

(a *Linear*) (*Nonlinear*)
 () [*Forecast Periods*](가 5),
 (가 3), (*Testing Periods*)(가 , 5)

Run
 (clipboard)

(layers) (input) (parameter)
 (calibrated) 가
 (layers)
 3 ()
 (calibration)

()

가 “ crisp
0 1 가
2 가 (weighting schema)

5.57 Risk Simulator

가

[Risk Simulator Forecasting | | Combinatorial Fuzzy Logic Risk Simulator | ROV BizStats | Combinatorial Fuzzy Logic Risk Simulator](#)
[ROV BizStats](#) 5.57

 Risk Simulator | Forecasting | Combinatorial Fuzzy Logic ()

 (가

 drop-list ()

(가 , 4,
12,) (가 , 5)

 Run

(clipboard)

가 Risk Simulator 가

Risk Simulator (forecasting methodologies)

..

5.56-

BizStats UI() Grid ,

.

RS , .

가 .

(layers)

(Testing Period)

,

, “ , 가

.

(paste)

2 : , , :

()

()

,

(Vari) (run) (cancel)

신경망 예측

STEP 1: 데이터 수동으로 데이터를 입력하고 다른 어플리케이션으로부터 붙이거나, 혹은 분석과 샘플 데이터를 시작할 수 있습니다

붙이기

N	VAR2	VAR3	VAR4	VAR5	VAR6	VAR7	VAR8	VAR9	VAR10	VAR11
NOT...		NNET								
1	1	459.11								
2	2	460.71								
3	3	460.34								
4	4	460.68								
5	5	460.83								
6	6	461.68								
7	7	461.66								
8	8	461.64								
9	9	465.97								
10	10	469.38								

STEP 2: 실행을 위한 분석 타입, 변수, 예측을 선택하십시오

한국어

VAR2

격지: 3

테스트 세트: 210

예측 기간: 210

복사

실행

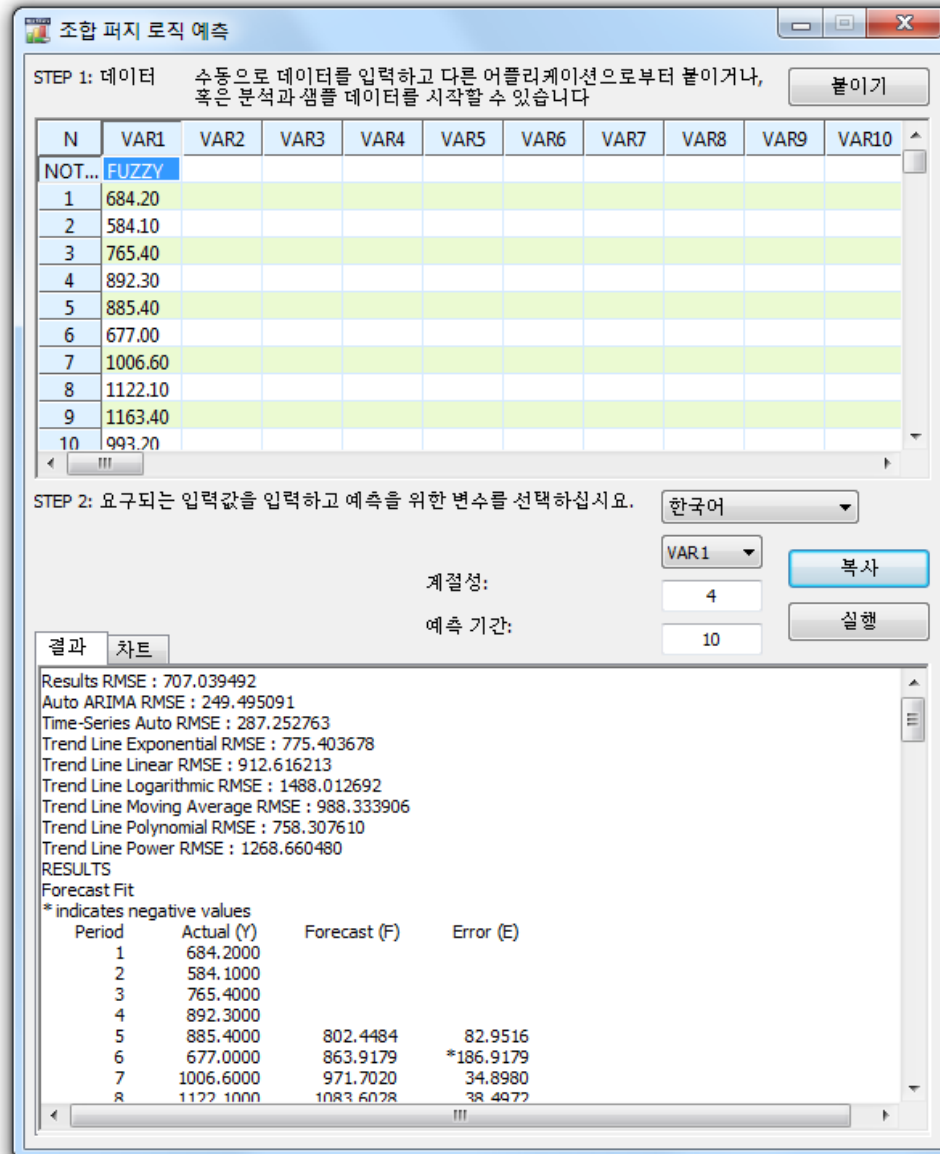
결과 차트 ☒ 멀티 단계 최적화 적용

Sum of Squared Errors (Training) : 1.822044
 RMSE (Training) : 0.093820
 Sum of Squared Errors (Modified) : 59375.218349
 RMSE (Modified) : 16.814849

Forecasting
 * indicates negative values

Period	Actual (Y)	Forecast (F)	Error (E)
211	581.5000	613.3528	*31.8528
212	584.2200	613.5197	*29.2997
213	589.7200	613.6203	*23.9003
214	590.5700	613.7188	*23.1488
215	588.4600	613.8520	*25.3920
216	586.3200	614.0608	*27.7408
217	591.7100	614.2046	*22.4946
218	593.2600	614.3029	*21.0429
219	592.7200	614.4223	*21.7023
220	592.3000	614.5671	*22.2671
221	589.2900	614.7154	*25.4254
222	593.9600	614.8963	*20.9363
223	597.3400	614.9954	*17.6554
224	600.0700	615.0992	*15.0292
225	596.8500	615.2115	*18.3615

Figure 5.56 – Neural Network Forecast()



5.57 -

BizStats UI() Grid ,

RS ,

(paste)

2 ;

, , , , (copy)

Goal Seek (Optimizer Goal Seek)

Goal Seek

, Risk Simulator

| Tools | Goal Seek

. Goal Seek

Risk Simulator s (advanced Optimization
routine) . 5.58 Goal Seek 가 가

5.58 (Goal Seek)

(One Variable Target Seek)

(Set cell) (Result)

(to value)

(By changing cell)

(run) (cancel)

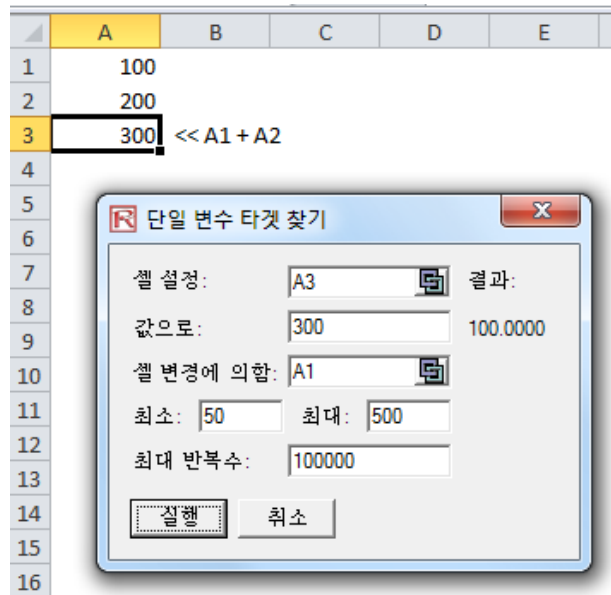


Figure 5.58 – Goal Seek

(Single Variable Optimizer /

)

•

Goal Seek

가

[Risk Simulator | Tools | Single Variable](#)

[Optimizer](#) (5.59)

가

Risk Simulator s

(advanced Optimization routines)

Risk Simulator

가 ,

5.59-

()

(One Variable Quick
Optimizer)
(ObjectiveCell) (Variable)
(Maximize) (Minimize)
(Tolerance) (Max Iteration)
(Min) (Max)
(Optimized Variable)
(Optimized Objective) (run) (cancel)

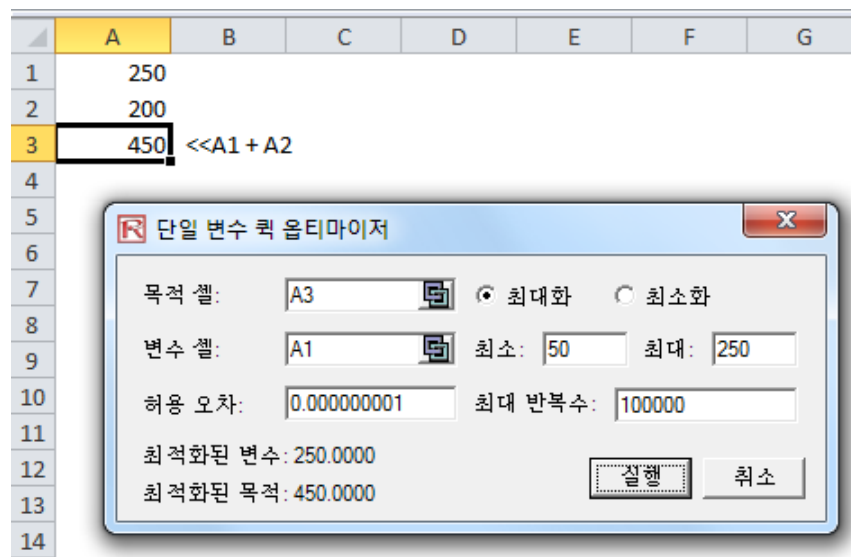


Figure 5.59 – Single Variable Optimizer ()

(Genetic Algorithm Optimization)

Risk Simulator | | (5.60)
(calibrating)
(가
), (Gradient Search
Test)가 . ((re-run)
(choice)
).
: (global)
()
. ()
).
(*Apply Gradient Search*
Test selection)(5.60) (re-run)
가 가
가
가 , Risk Simulator's Optimization
module , Risk-based dynamic

5.60- (Genetic Algorithm)
(Objective Cell)

가(Add)
(Column /Cell) (ColumnMin) (ColumnMax)
(Constraints)
(Max Iterations) (Mutation Rate)
(Population Size) (Diversity)

(Crossover Rate) (Elitism)
 (Unchange) (run) (Result)

유전적 알고리즘

목적 셀: ☒ 최대화 ☐ 최소화

변수:

셀	최소	최대

제약조건:

셀	최소	최대

최대 반복수: 돌연변이율:

집단 크기: 다양성:

크로스오버 비율: 엘리트주의:

크로스오버: 불변하는:

☐ 증감률 검색 테스트 적용

결과:

Figure 5.60 –

ROV

Decision Tree

ROV Decision Tree (5.61) 결정 트리는 결정 트리 모델을 만들고 가치를 평가하기 위해 사용됩니다. 다음과 같은 고급 방법론 및 분석기법들이 추가로 포함되어 있습니다.

- 결정 트리 모델
- 몬테카를로 리스크 시뮬레이션
- 민감도 분석
- 시나리오 분석
- 베이지안 방법 (결합 및 사후 확률 갱신)
- 정보에 대한 기대 값
- MINIMAX
- MAXIMIN
- 리스크 프로파일

다음 사항들은 이 손쉬운 도구를 사용하기 위해 필요한 몇 가지 주요한 빠른 사용법 습득 요령과 절차들입니다.

- 본 모듈은 11 가지의 언어로 사용 가능하며, 언어 메뉴를 통해 현재 사용되는 언어를 변경할 수 있습니다..
- *옵션* 노드나 *터미널* 노드들을 삽입하려면, 먼저 기존의 노드를 선택하고 나서 *옵션* 노드 아이콘 (사각형)이나 *터미널* 노드 아이콘 (삼각형)을 클릭하거나, 아니면 *삽입* 메뉴에 있는 기능을 사용합니다.
- 개별 *옵션* 노드나 *터미널* 노드의 속성들을 수정하려면, 노드를 더블 클릭합니다. 노드를 클릭하면, 때로는 그 아래에 있는 하부 노드들 역시 모두 선택됩니다 (이 기능은 선택된 노드로부터 전체 트리 시작을 가능하게 합니다). 오직 특정 노드만 선택하고 싶은 경우, 빈 배경 위를 클릭한 후 해당 노드를 다시 클릭하면, 해당 노드만 개별적으로 선택됩니다. 또한, 현 설정 상태에 따라 (우측 클릭을 하여 또는 *편집* 메뉴에서 *개별 노드 이동* 아니면 *단체 노드 이동*을 선택합니다) 개별 노드나 선택된 노드로부터 시작된 전체 트리를 이동할 수도 있습니다.
- 다음은 노드 속성 사용자 인터페이스에서 사용자가 정의하고 구성할 수 있는 것들에 대한 일부 개략적인 설명들입니다. 전략 트리 내 어떤 영향을 주는지 알아보려면 다음 각 항목 별로 다른 설정을 시도해 보는 것이 가장 간편한 방법입니다.
 - o *이름*. 이름은 노드 위에 나타납니다.
 - o *값*. 값은 노드 아래 나타납니다.
 - o *엑셀 링크*. 엑셀 스프레드시트의 셀에 있는 값을 연결합니다.
 - o *노트*. 노트는 노드의 위나 아래에 삽입할 수 있습니다.
 - o *모델로 보기*. 이름, 값, 노트들의 조합을 보여줍니다.

- o **구역 색/전체 색.** 노드 색상은 한 노드에만 국한해서, 또는 전체적으로 변경할 수 있습니다.
 - o **라벨 내부 모양.** 노드 안에 문구를 넣을 수 있습니다 (문구가 길어질 경우, 노드 폭을 조정해야 할 수도 있습니다.).
 - o **브랜치 이벤트 이름.** 노드를 유발시킨 이벤트를 표시하기 위해 그 노드로 이어지는 브랜치상에 문구를 넣을 수 있습니다.
 - o **실 옵션 선택.** 현 노드에 특정 실 옵션 형태를 배정할 수 있습니다. 노드에 실 옵션을 배정함으로써, 필요한 입력 변수들의 목록 생성이 가능합니다.
- **옵션 노드, 터미널 노드, 문자 상자들의 요소들을 포함하여 전체 요소들은 모두 사용자 정의가 가능합니다.** 예를 들어, 각 요소들에 대해 다음과 같은 설정들을 변경할 수 있습니다.
 - o 이름, 값, 노트, 라벨, 이벤트 이름의 글꼴 설정.
 - o 노드 크기(최소 및 최대 높이와 폭).
 - o 경계선(선 스타일, 폭, 색).
 - o 그림자(색, 그림자 적용 여부).
 - o 전체 색.
 - o 전체 모양.
- 편집 메뉴의 **데이터 요구사항 윈도우 보기**를 선택하면 전략 트리의 우측 에 창이 열리며, 옵션 노드나 터미널 노드가 선택되었을 때 그 노드의 속성들을 보여주고 수정도 바로 가능합니다. 매번 노드 위 에서 더블 클릭을 하는 대신 이 기능을 사용할 수 있습니다.
- 전략 트리를 구축하기 시작하는데 도움이 되도록, **파일** 메뉴 안에 **사례 파일들**을 사용할 수 있습니다.
- **파일** 메뉴에서 **파일 보호**를 선택하면 최대 256-비트의 비밀번호로 전략 트리를 암호화 할 수 있습니다. 비밀번호를 잃어버렸을 경우 파일을 더 이상 열 수 없으므로 파일을 암호화 할 경우 주의가 필요합니다.
- **화면 캡처** 또는 기존 모델 인체는 **파일** 메뉴를 통해 가능합니다. 캡처한 화면은 다른 소프트웨어 애플리케이션으로 붙여 넣기가 가능합니다.
- **전략 트리 추가, 복사, 이름 바꾸기, 삭제**는 전략 트리 탭에서 우측 클릭을 하거나 **편집** 메뉴를 통해 가능합니다.
- 어느 옵션이나 터미널 노드 상에서도 **파일 링크 삽입**과 **코멘트 삽입**이 가능하며, 배경이나 캔버스 부위 어느 곳이나 문구나 그림을 삽입할 수도 있습니다.
- **기존 스타일**을 변경하거나 또는 전략 트리의 사용자 스타일을 만들고 관리할 수 있습니다 (이 기능은 전체 전략 트리의 크기, 모양, 계획, 글꼴 크기/색 사양들을 포함합니다).
- **결정**이나 **불확실성** 또는 **터미널** 노드들을 삽입하려면, 먼저 기존의 노드를 선택하고 나서 결정 노드 아이콘 (사각형)이나 불확실성 노드 아이콘 (원형) 또는 터미널 노드 아이콘 (삼각형)을 클릭하거나, 아니면 **삽입** 메뉴에 있는 기능들을 사용합니다.

- 개별 결정이나 불확실성 또는 터미널 노드의 속성들을 수정하려면, 노드를 더블 클릭합니다. 다음은 노드 속성 사용자 인터페이스에서 사용자가 정의하고 구성할 수 있는 결정 트리 모듈의 추가적인 고유 항목들의 일부입니다
 - o 결정 노드: 노드 상의 값을 자동 계산 하거나 사용자가 그 값을 고칠 수 있습니다. 자동 계산 옵션이 기본적으로 설정되며, 완성된 결정 트리 모델 상에서 실행을 클릭하면, 결정 노드들이 자동 계산된 결과들로 수정됩니다.
 - o 불확실성 노드: 이벤트 명, 확률, 시뮬레이션 가정 설정. 불확실성 브랜치가 만들어지고 난 후에만, 확률 이벤트 명, 확률, 시뮬레이션 가정들을 추가할 수 있습니다.
 - o 터미널 노드: 수동 입력, 엑셀 링크, 시뮬레이션 가정 설정. 터미널 이벤트 청산을 수동으로 입력하거나, 엑셀의 셀과 연결하거나 (예, 이 청산을 계산하는 대형 엑셀 모델을 가지고 있다면, 모델을 이 엑셀 모델의 출력 셀과 연결할 수 있습니다), 실행 시뮬레이션에 대한 확률 분포 가정들을 설정할 수 있습니다.
- 편집 메뉴에서 노드 속성 윈도우 보기를 사용할 수 있으며, 노드를 선택하면 선택된 노드의 속성들이 새로 고쳐집니다.
- 결정 트리 모듈 역시 다음과 같은 분석기법들을 가지고 있습니다.
 - o 결정 트리에 대한 몬테카를로 시뮬레이션 모델링
 - o 사후 확률들을 얻기 위한 베이지 분석
 - o 완벽한 정보에 대한 기대 값, MINIMAX 및 MAXIMIN 분석, 리스크 프로파일, 불완전한 정보의 가치
 - o 민감도 분석
 - o 시나리오 분석
 - o 효용 함수 분석

Simulation Modeling

본 도구는 결정 트리에 대해 몬테카를로 리스크 시뮬레이션을 실행합니다 (그림 5.62). 이 도구를 사용하여 시뮬레이션 실행을 위한 입력 가정으로 확률 분포를 설정할 수 있습니다. 선택한 노드에 대한 가정을 설정하거나, 또는 새 가정을 설정하여 수치 방정식이나 수식에서 이러한 새 가정을 사용할 수 있습니다(또는 이전에 생성한 가정 사용). 예를 들어, Normal(예: 평균이 100 이고 표준 편차가 10 인 정규 분포)이라는 새 가정을 설정하고 의사 결정 트리에서 시뮬레이션을 실행하거나 $(100 * \text{Normal} + 15.25)$ 와 같은 수식에서 이 가정을 사용할 수 있습니다. 수식 상자에서 자신의 모델을 생성하십시오. 기본 계산을 사용하거나, 기존 변수 목록을 두 번 클릭하여 기존 변수를 수식에 추가할 수 있습니다. 새 변수는 필요에 따라 수식이나 가정으로 목록에 추가할 수 있습니다.

Bayes Analysis

이 베이지안 분석 도구는 하나의 경로를 따라 연결된 2 개의 불확실성 사건에 대해 수행할 수 있습니다 (그림 5.63). 예를 들어, 오른쪽 예제에서 불확실성 A 와 B 는 연결되어 있으며, 이 때 먼저 사건 A 가 발생한 다음 사건 B 가 발생합니다. 첫 번째 사건 A 는 2 개의 결과(유리 또는 불리)가 있는 시장 조사입니다. 두 번째 사건 B 도 2 개의 결과(강함 및 약함)를 가지는 시장 조건입니다. 이 도구는 사전 확률과 신뢰성 조건부 확률을 입력하여 결합, 경계 및 베이지안 사후 갱신 확률을 계산하는 데 사용합니다. 또는 신뢰성 확률은 사후 갱신 조건부 확률이 있을 때 계산할 수 있습니다. 아래에서 원하는 관련도 분석을 선택하고 예제 로드를 클릭하면 선택한 분석에 해당하는 샘플 입력과 오른쪽 모눈에 표시되는 결과뿐만 아니라 그림의 의사 결정 트리에서 입력으로 사용되는 결과가 무엇인지 확인할 수 있습니다.

Quick Procedures

- 단계 1: 첫 번째 및 두 번째 불확실성 사건에 대한 이름을 입력하고 각 사건이 가지는 확률 사건(자연 상태 또는 결과) 개수를 선택합니다.
- 단계 2: 각 확률 사건 또는 결과의 이름을 입력합니다.
- 단계 3: 두 번째 사건의 사전 확률 및 각 사건이나 결과에 대한 조건부 확률을 입력합니다. 확률의 합계는 100%가 되어야 합니다.

Expected Value of Perfect Information, MINIMAX and MAXIMIN Analysis, Risk Profiles, and Value of Imperfect Information

이 도구는 완전 정보의 기대 가치(EVPI), 최소최대 및 최대최소 분석뿐만 아니라 위험 프로필과 불완전 정보의 가치를 계산합니다 (그림 5.64). 시작하려면, 고려 중인(예: 대규모, 중규모, 소규모 설비 구축) 의사 결정 분기 또는 전략의 수와 불확실한 사건 또는 자연 결과 상태(예: 유리한 시장, 불리한 시장)의 수를 입력하고 각 시나리오 아래에 기대 이익을 입력합니다.

완전 정보의 기대 가치(EVPI)는 수요를 미리 정확히 알고 있다는 가정 하에(시장 조사 또는 확률적 결과를 더 잘 식별할 수 있는 기타 수단을 통해) 자연의 확률적 상태에 대한 나이브 추정값과 비교해, 이러한 정보에 추가된 값이 있는지 여부를 계산합니다(예: 시장 조사에서 값이 추가되는 경우). 시작하려면, 고려 중인(예: 대규모, 중규모, 소규모 설비 구축) 의사 결정 분기 또는 전략의 수와 불확실한 사건 또는 자연 결과 상태(예: 유리한 시장, 불리한 시장)의 수를 입력하고 각 시나리오 아래에 기대 이익을 입력합니다.

최소최대(최대 후회의 최소화) 및 최대최소(최소의 후회 최대화)는 최적의 의사 결정 경로를 찾기 위한 2 개의 대체적 접근 방법입니다. 이 두 개의 접근법은 잘 사용되지 않지만, 여전히 의사 결정 과정에 필요한 통찰력을 제공해 줍니다. 존재하는(예: 대규모, 중규모 또는 소규모 설비 구축) 의사 결정 분기 또는 경로 개수뿐만 아니라 각 경로 아래에 불확실성 사건이나 자연 상태(예: 유리한 경제, 불리한 경제)의 수를 입력합니다. 그리고

여러 시나리오의 수익 테이블을 완료하고 최소최대 및 최대최소 결과를 계산합니다. 예제 로드를 클릭해도 샘플 계산을 확인할 수 있습니다.

Sensitivity

의사 결정 경로의 값에 미치는 영향을 확인하기 위해 입력 확률의 민감도 분석을 수행합니다 (그림 5.65). 먼저, 아래에서 분석할 의사 결정 노드를 하나 선택하고, 목록에서 테스트할 확률 사건을 하나 선택합니다. 동일한 확률의 불확실성 사건이 여러 개 있는 경우 개별적으로 또는 동시에 분석할 수 있습니다.

민감도 차트는 다양한 확률 수준 아래에 의사 결정 경로의 값을 표시합니다. 숫자값은 결과 테이블에 표시됩니다. 십자선이 있는 경우 십자선의 위치는 특정 의사 결정 경로가 다른 경로보다 우세한 확률 사건을 나타냅니다.

Scenario Tables

입력이 일부 변경된 경우, 시나리오 테이블을 생성하여 출력값을 확인할 수 있습니다 (그림 5.66). 시나리오 테이블의 입력 변수로는 분석할 하나 이상의 의사 결정 경로(선택된 각 경로의 결과는 개별 테이블 및 차트로 표시됨), 그리고 하나 또는 두 개의 불확실성 노드나 만기 노드를 선택할 수 있습니다.

Quick Procedures

- 아래의 목록에서 분석할 의사 결정 경로를 하나 이상 선택합니다
- 사건의 확률을 개별적으로 변경할 것인지 또는 모든 동일한 확률 사건을 한 번에 변경할 것인지 결정합니다.
- 입력 시나리오 범위 입력.

Utility Function Generation

효용 함수인 (그림 5.67) $U(x)$ 는 때때로 의사 결정 트리에서 만기 수익의 기대 가치 대신 사용됩니다. $U(x)$ 는 모든 가능한 결과에 대한 지루하고 자세한 실험법을 사용하거나 지수 외삽법(여기에서 사용)을 사용하여 개발할 수 있습니다. 이러한 방법은 위험 회피형 의사 결정자(하강 시 잠재력이 동일한 상승보다 피해가 더 큼), 위험 중립형 의사 결정자(상승과 하강이 동일한 매력을 가짐) 또는 위험 선호형 의사 결정자(상승 잠재력이 더 매력적임)에 대해 모델링할 수 있습니다. 효용 곡선 및 테이블을 계산하려면 만기 수익의 최소/최대 기대 가치와 그 사이에 있는 데이터 지점 수를 입력해야 합니다.

참여할 경우 $\$X$ 를 벌거나 $\$X/2$ 를 잃을 수 있고, 참여하지 않으면 $\$0$ 의 수익을 내는 50:50 게임을 할 경우, 이 $\$X$ 는 얼마입니까? 예를 들어, 전혀 참여하지 않을 때와 확률이 동일한 경우 $\$100$ 를 벌거나 $\$50$ 를 잃을 수 있는 베팅 사이에서 위험을 고려하지 않는다면 X 는 $\$100$ 입니다. 아래에서 양의 소득 상자에 X 를 입력합니다. X 가 클수록 위험을 덜 꺼려하고, 이에 반해 X 가 작을수록 보다 더 위험을 꺼려한다는 점을 참고하십시오.

필요한 입력을 입력하고, $U(x)$ 유형을 선택하고, 계산 유틸리티를 클릭하여 결과를 획득합니다. 또한 계산된 $U(x)$ 값을 의사 결정 트리에 적용하여 재실행하거나, 수익의 기대 가치를 사용하기 위해 트리를 되돌릴 수 있습니다.

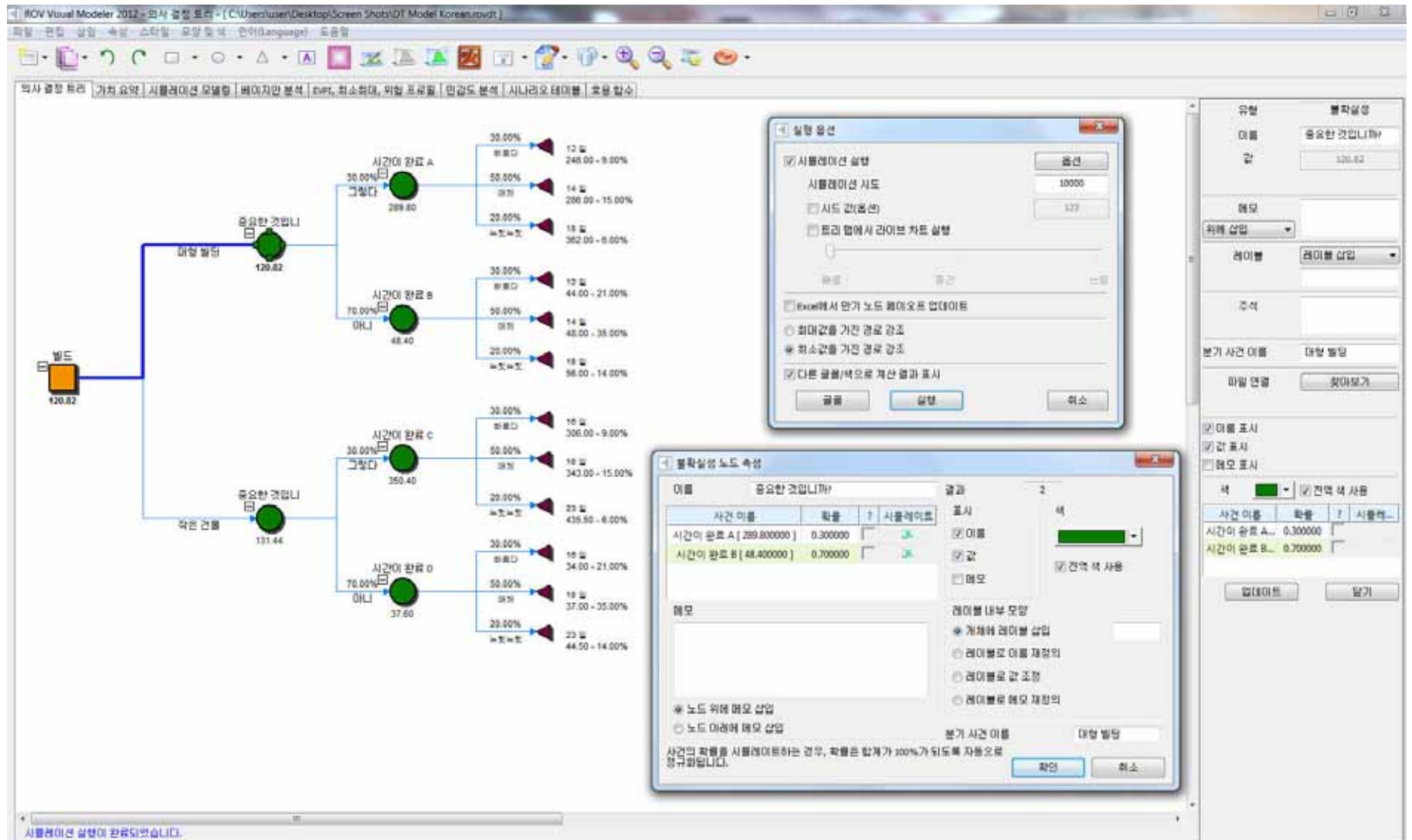


그림 5.61 – ROV 결정 트리 (결정 트리)

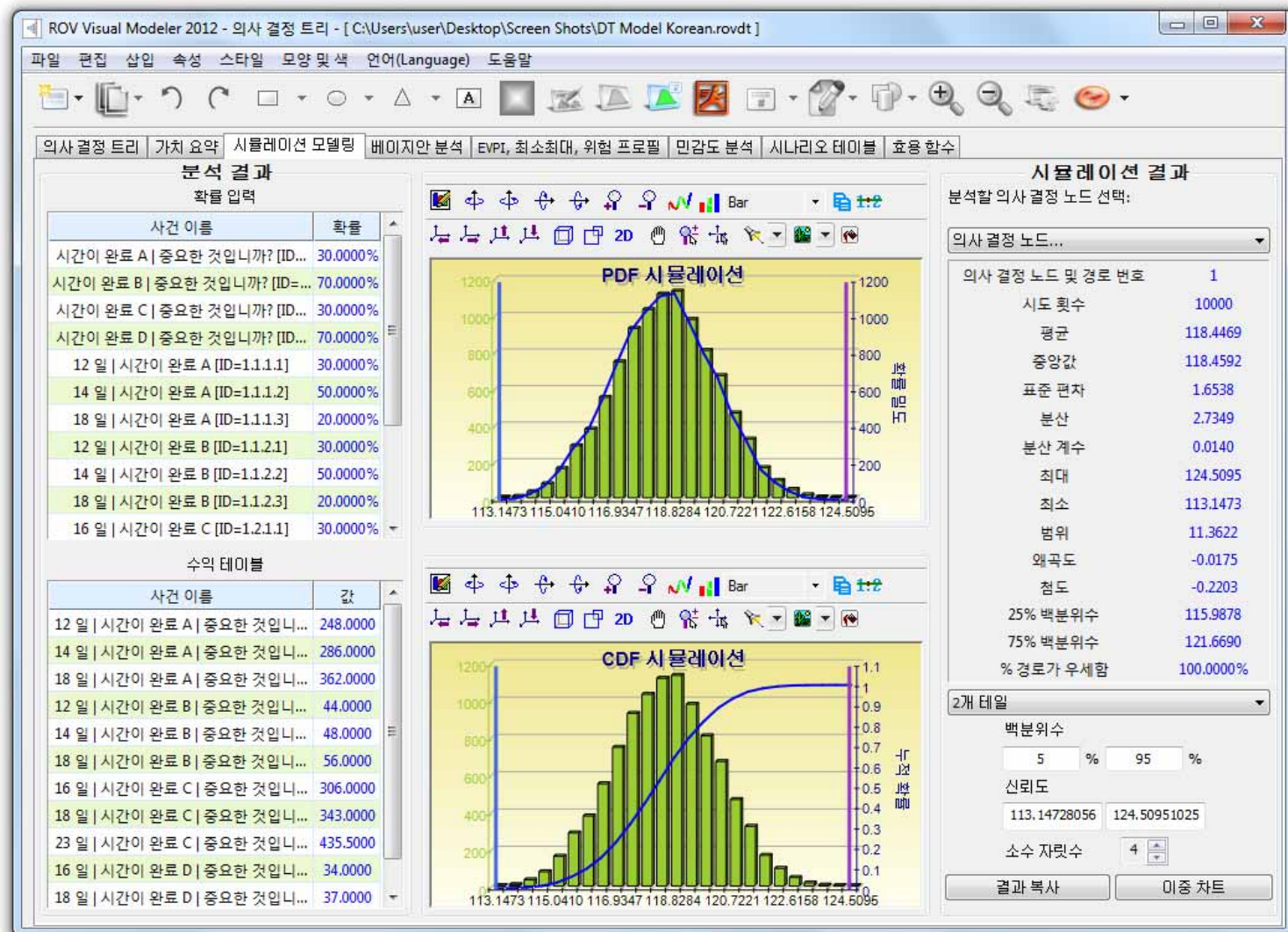


그림 5.62 – ROV 결정 트리 (시뮬레이션 결과)

ROV Visual Modeler 2012 - 의사 결정 트리 - [C:\Users\user\Desktop\Screen Shots\DT Model Korean.rovdt]

파일 편집 삽입 속성 스타일 모양 및 색 언어(Language) 도움말

의사 결정 트리 | 가치 요약 | 시뮬레이션 모델링 | 베이직 분석 | EVPI, 최소최대, 위험 프로필 | 민감도 분석 | 시나리오 테이블 | 효용 함수

이 베이직 분석 도구는 하나의 경로를 따라 연결된 2개의 불확실성 사건에 대해 수행할 수 있습니다. 예를 들어, 오른쪽 예제에서 불확실성 A와 B는 연결되어 있으며, 이때 먼저 사건 A가 발생한다. 다음 사건 B가 발생합니다. 첫 번째 사건 A는 2개의 결과(유리 또는 불리)가 있는 시장 조사입니다. 두 번째 사건 B도 2개의 결과(강한 및 약한)를 가지는 시장 조건입니다. 이 도구는 사전 확률과 신뢰성 조건부 확률을 입력하여 결합, 경계 및 베이직의 사후 갱신 확률을 계산하는 데 사용됩니다. 또는 신뢰성 확률은 사후 갱신 조건부 확률이 있을 때 계산할 수 있습니다. 아래에서 원하는 관련도 분석을 선택하고 예제 로드를 클릭하면 선택한 분석에 해당하는 샘플 입력과 오른쪽 모눈에 표시되는 결과뿐만 아니라 그림의 의사 결정 트리에서 입력으로 사용되는 결과가 무엇인지 확인할 수 있습니다.

☒ 사전 및 신뢰성 결합 확률을 고려하여 베이직 갱신 사후 확률을 계산합니다(보다 일반적).

☐ 사전 및 사후 확률을 고려하여 신뢰성 결합 확률을 계산합니다(덜 일반적).

단계 1: 첫 번째 및 두 번째 불확실성 사건에 대한 이름을 입력하고 각 사건이 가지는 확률 사건(자연 상태 또는 결과) 개수를 선택합니다.

첫 번째 사건 이름: 확률 사건 또는 상태:

두 번째 사건 이름: 확률 사건 또는 상태:

단계 2: 각 확률 사건 또는 결과의 이름을 입력합니다.

상태	Market Research	Market Conditions
1	Favorable	Strong
2	Unfavorable	Weak

예제 로드

계산

저장된 모델 이름:

추가

삭제

베이직 분석 결과

사전 확률 및 신뢰성 조건부 확률

확률 (Strong)	확률 (Weak)
45.00%	55.00%
확률 (Favorable Strong)	확률 (Favorable Weak)
60.00%	30.00%
확률 (Unfavorable Strong)	확률 (Unfavorable Weak)
40.00%	70.00%

결합 및 경계 확률

확률 (Favorable)	확률 (Unfavorable)
43.50%	56.50%
확률 (Strong $\cap$ Favorable)	확률 (Weak $\cap$ Favorable)
27.00%	16.50%
확률 (Strong $\cap$ Unfavorable)	확률 (Weak $\cap$ Unfavorable)
18.00%	38.50%

사후 또는 갱신 확률

확률 (Strong Favorable)	확률 (Weak Favorable)
62.07%	37.93%
확률 (Strong Unfavorable)	확률 (Weak Unfavorable)
31.86%	68.14%

Market Research Reliability:
60% (Favorable given Strong)
40% (Unfavorable given Strong)
30% (Favorable given Weak)
70% (Unfavorable given Weak)

Second Uncertainty (F)
62.07% Strong Market
37.93% Weak Market
PRIOR PROBABILITIES: 45% AND 55%

Second Uncertainty (U)
31.86% Strong Market
68.14% Weak Market
PRIOR PROBABILITIES: 45% AND 55%

그림 5.63 – ROV 결정 트리 (베이즈 분석)

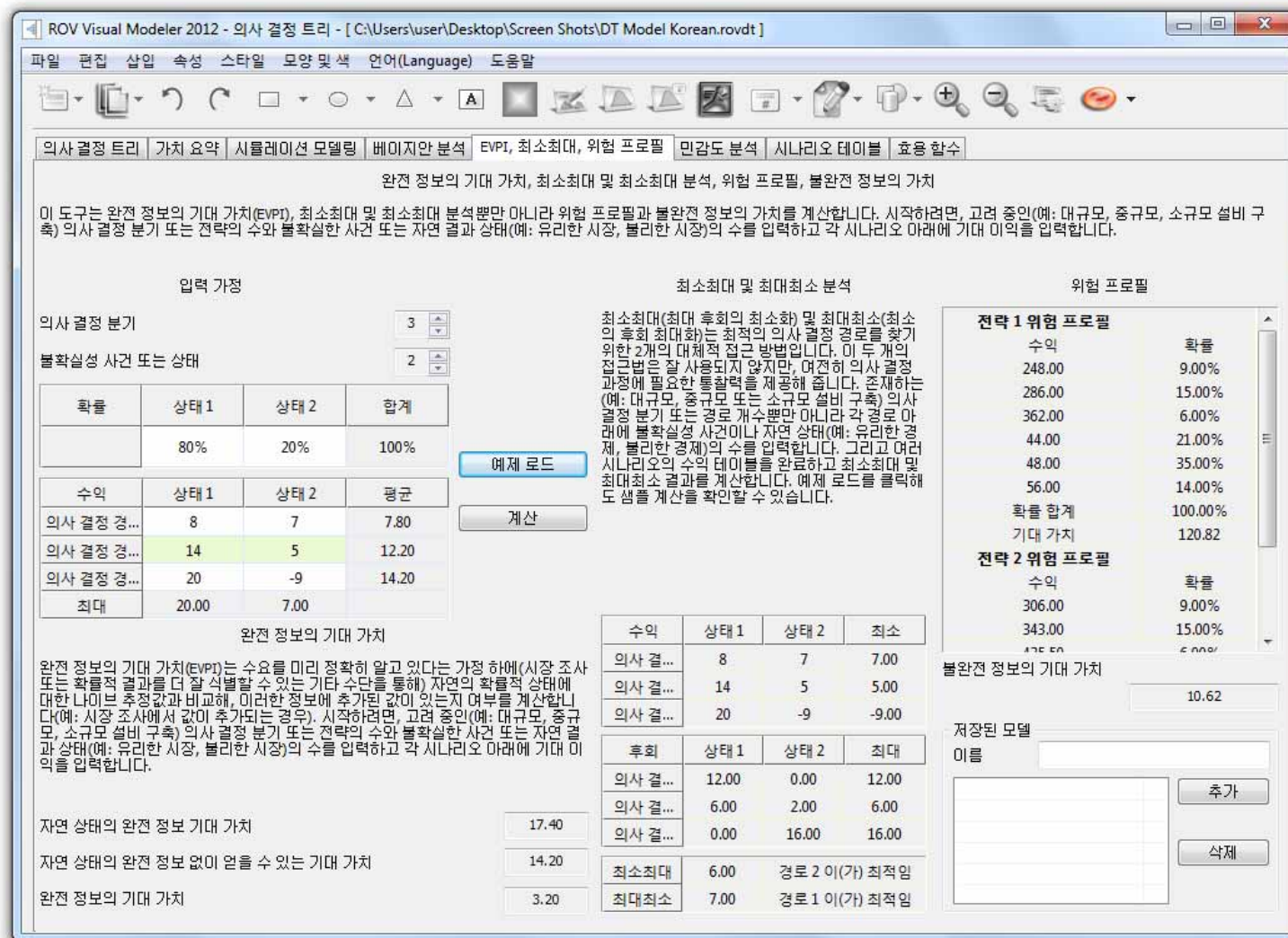


그림 5.64 – ROV 결정 트리 (EVPI, MINIMAX, 리스크 프로파일)

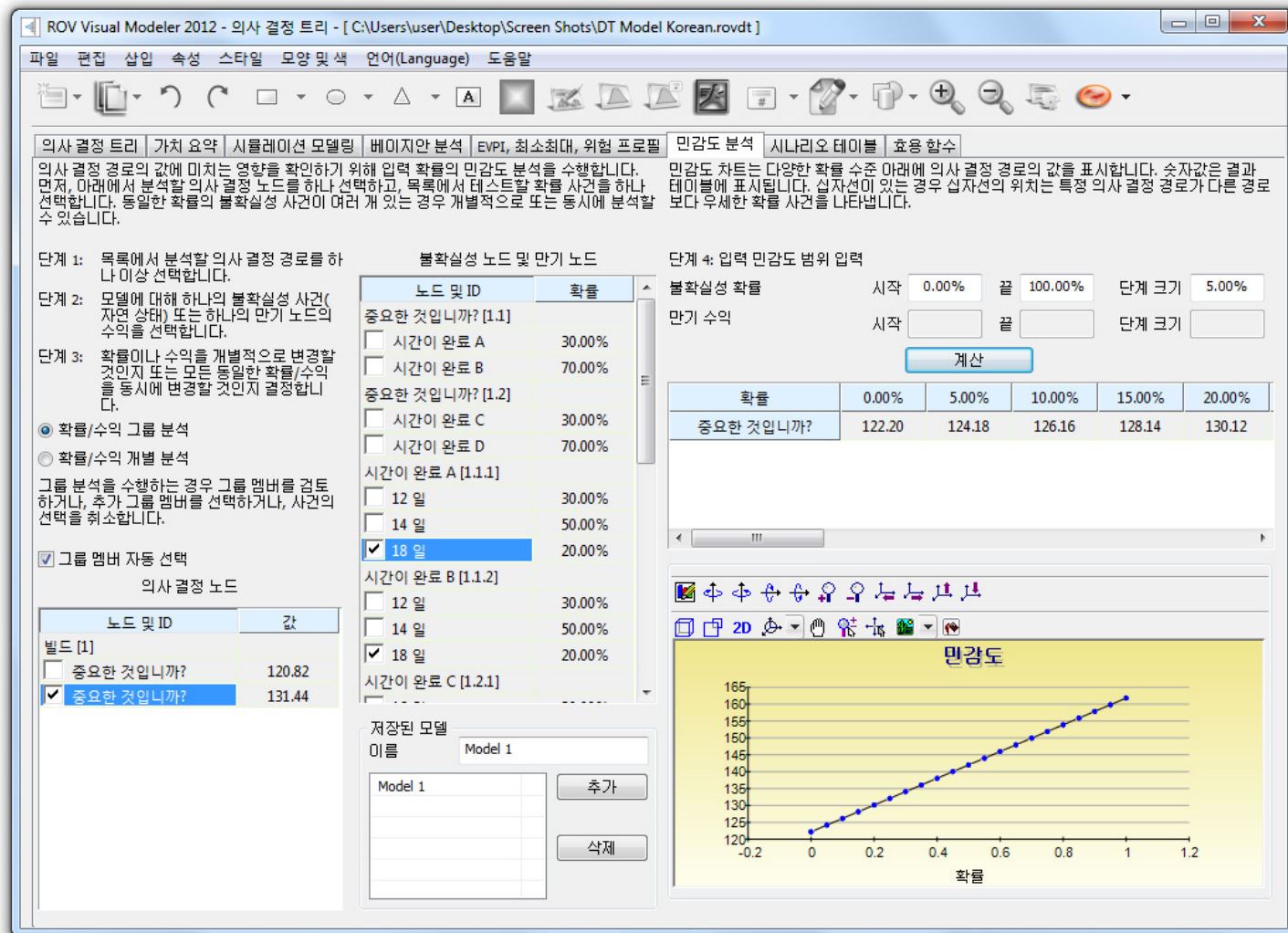


그림 5.65 – ROV 결정 트리 (민감도 분석)

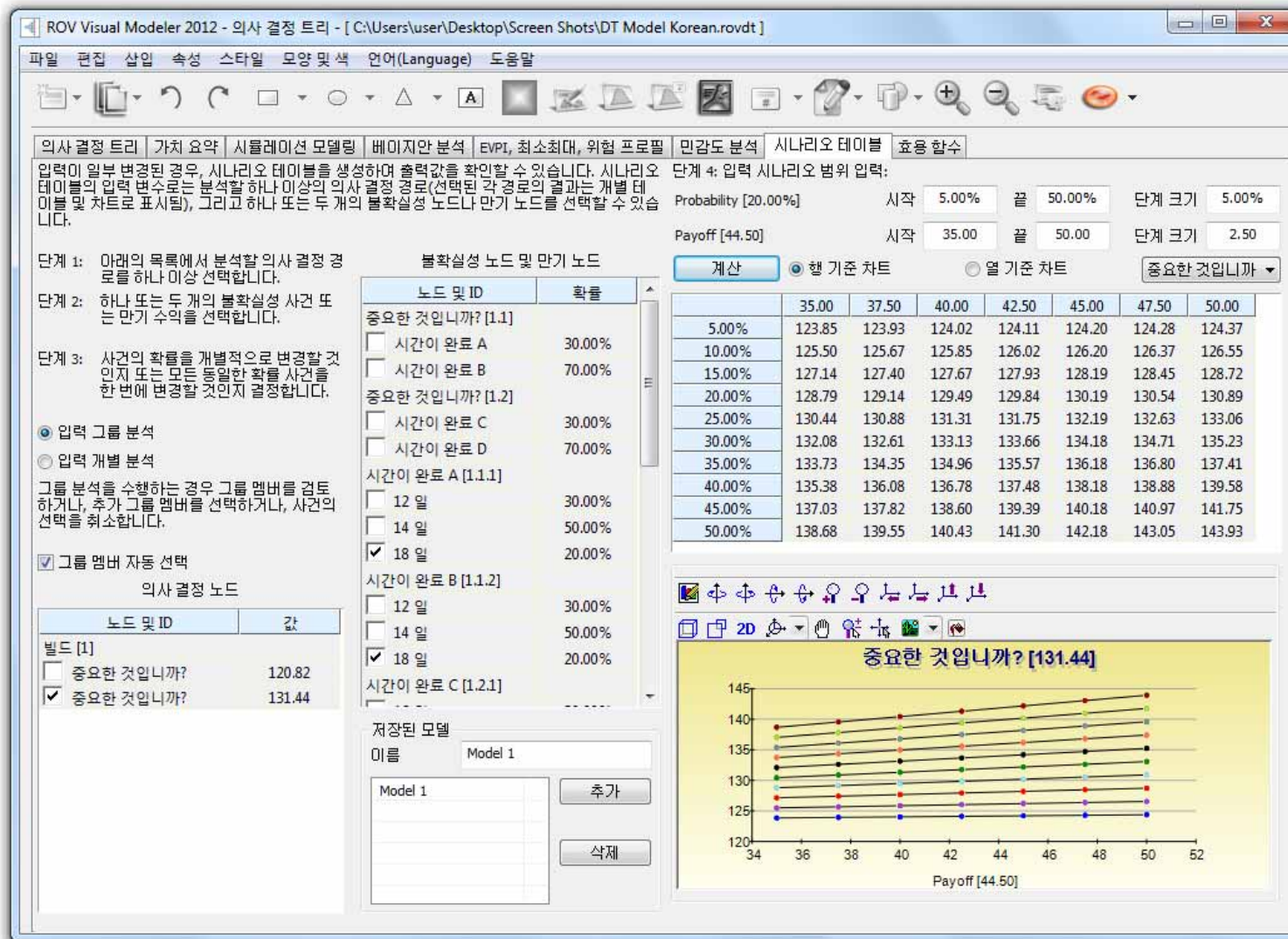


그림 5.66 – ROV 결정 트리 (시나리오 표)

: 가 (가)

- — , (: Normal

W Weibull).

- — , (large icons, small icons, list).

- — (:

가 가

- — (가).

, 가

- 가 — 가 가 가 가

가

(,).

: /

- / — Risk Simulator

, , , , , , , 가 , ,

Risk Simulator

, 가 .

Risk Simulator

, , , 가

Risk Simulator

, Risk Simulator 가 , ,

, , , ,

.

Risk Simulator

.

- / — / (가)

:

- 가 — 가 ()

- — ()

- — (, ,)

:

- — Statistical Analysis and Data Diagnostic , , mean-reversion, jump-diffusion

.

.

,

(: Brownian Motion, Mean-Reversion, Jump-Diffusion,) 가

가 . 가 가

(:

Brownian Motion 가 가
가
,).
가 ,
가
가
가 (: mean-reversion is 110% mean-reversion
가).

- :
- — Risk Simulator 42 가 PDF, CDF, ICDF ,
 - —
(: [2, 2], [3, 5], [3.5, 8] 가
PDF, CDF, ICDF ,).
 - — (가
),

- :
- — .
 - —
().
- 가

- — 가 view
- — 가 view
- —normal view global view

- — 가
- RMSE—

- : **ARIMA**
- — (: 가 , 100 가 105 가 , ARIMA

- : **Basic Econometrics**
- —

: : *Logit, Probit, and Tobit*

- — logit probit (0, 1)
, tobit .

: : *Stochastic Processes*

- Default Sample Inputs— ,
- Statistical Analysis Tool for Parameter Estimation —
- Stochastic Process Model — 가
, 가 (:
mean-reversion is 110% mean-reversion 가
).

: : *Trendlines*

- — 가 .
- RS Functions — set input
assumption get forecast statistics .
RS function (Start > Programs > Real Options
Valuation > Risk Simulator > Tools > Install Functions), RS
function .

- :
 - — 가
가 Start > Programs > Real Options Valuation > Risk Simulator shortcut location .
 - — 가
(www.realoptionsvaluation.com/download.html
www.rovdownloads.com/download.html).
 - : **ID**
 - HWID — Install License user
interface HWID ,
HWID
E-mail HWID link .
 - —Start, Programs, Real Options Valuation, Risk Simulator
Troubleshooter , HWID Get HWID
- : **(LHS)** **(MCS)**
 - — 가 , Risk Simulator, Options
menu Monte Carlo setting . Latin Hypercube
Sampling correlated copula method .
 - LHS Bins — bin
 - —Options
- :
 - , , , —

(www.realoptionsvaluation.com/download.html
www.rovdownloads.com/download.html).

- :
 가 — 가 가 ,
 (=) (>= <=)
- :
 —
- :
 — , 가 , ,
 ,
 가 . 가 가
- :
 —
- :
 — ,
- :
 — 가
- :
 , 가
- :
 —
- :
 — (가 , , , ,)

가

• Risk Simulator 가

:

• (Risk Simulator , 가 , ,)

•

(save open).

• 가

• Risk Simulator Risk Simulator *.RiskSim

• 가

:

• 가 (ROV Risk Simulator Advanced Subtractive Random Shuffle 가

가

: (SDK) DLL

- SDK, DLL, OEM—Risk Simulator

가

DLL

admin@realoptionsvaluation.com

: **Risk Simulator**

- ROV Troubleshooter—

HWID

, Risk Simulator가

가

가

troubleshooter

- Risk Simulator — Risk Simulator가 Start, Programs, Real Options Valuation, Risk Simulator shortcut location

Risk Simulator, Options menu

:

- — , 가

가

- —

: **Tornado Analysis**

- Tornado Analysis —Tornado Analysis

Tornado
 precedent가
 (Show All Variables).
 () Tornado
 , precedent가
 (: Tornado 가
 5 200
),
 Tornado Analysis (:
 5 가 , Show Top 10 Variables .
 ,
 Tornado ; Tornado
).

- — $\pm 10\%$
 가 (Spider chart Tornado
 precedent).
- Zero Values and Integers—Tornado inputs with zero
 or integer values only .
 가 (:
 , = 1, = 2, = 3, , 1 +/- 10%
 0.9 1.1 , 가).
- — 가

: Troubleshooter

- ROV Troubleshooter — HWID ,
 , Risk Simulator가 가
 가 troubleshooter .