



RISIKO SIMULATOR

Benutzerhandbuch

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

Risk Simulator 2012

Dieses Handbuch und die darin beschriebene Software werden unter Lizenz zur Verfügung gestellt. Sie dürfen nur in Übereinstimmung mit den Bestimmungen des Endbenutzer-Lizenzvertrags verwendet und kopiert werden. Die Informationen in diesem Dokument werden lediglich zur Information bereitgestellt, und Veränderungen sind vorbehalten. Diese Informationen stellen keine Verpflichtung seitens Real Options Valuation, Inc, bezüglich der Handelsüblichkeit oder der Eignung für einen bestimmten Zweck dar.

Kein Teil dieses Handbuchs darf auf keine Art und Weise, elektronisch oder mechanisch einschließlich Fotokopierung und Aufzeichnung, ohne die ausdrückliche schriftliche Genehmigung von Real Options Valuation, Inc vervielfältigt oder übertragen werden.

Die Inhalte basieren auf urheberrechtlich geschützten Publikationen von Dr. Johnathan Mun, Gründer und Geschäftsführer, Real Options Valuation, Inc.

Verfasst von Dr. Johnathan Mun.

Verfasst, entwickelt und herausgegeben in den Vereinigten Staaten von Amerika.

Um zusätzliche Kopien dieses Dokuments zu erwerben, kontaktieren Sie Real Options Valuation, Inc. unter der nachfolgende E-Mail-Adresse:

Admin@RealOptionsValuation.com oder besuchen Sie www.realoptionsvaluation.com

© 2005-2012 von Dr. Johnathan Mun. All Rechte vorbehalten.

Microsoft® ist ein in der US und anderen Ländern eingetragenes Warenzeichen von Microsoft Corporation.

Andere in diesem Handbuch erwähnten Produktnamen können Warenzeichen und/oder eingetragene Warenzeichen der jeweiligen Inhaber sein.

INHALTSVERZEICHNIS

1. EINLEITUNG.....	8
WAS GIBT'S NEUES IN VERSION 2011/2012.....	13
<i>Ein detailliertes Leistungsverzeichnis des Risiko Simulators</i>	<i>13</i>
<i>Allgemeine Einsatzmöglichkeiten</i>	<i>13</i>
<i>Simulationsmodell.....</i>	<i>14</i>
<i>Prognose Modul.....</i>	<i>15</i>
<i>Optimierungsmodell.....</i>	<i>15</i>
<i>Analytisches Hilfsprogramm Modul</i>	<i>16</i>
<i>Statistiken und BizStats Modul Statistics</i>	<i>17</i>
2. MONTE-CARLO-SIMULATION	19
<i>Was ist die Monte-Carlo-Simulation?</i>	<i>19</i>
<i>Erste Schritte mit Risiko Simulator.....</i>	<i>20</i>
<i>Ein Überblick der Software</i>	<i>20</i>
<i>Eine Monte-Carlo-Simulation ausführen</i>	<i>21</i>
<i>1. Ein neues Simulationsprofil starten.....</i>	<i>21</i>
<i>2. Inputhypothesen definieren.....</i>	<i>24</i>
<i>3. Outputvorausberechnungen definieren</i>	<i>27</i>
<i>4. Die Simulation ausführen</i>	<i>28</i>
<i>5. Die Vorausberechnungsergebnisse interpretieren</i>	<i>29</i>
<i>Korrelationen und Präzisionskontrolle.....</i>	<i>36</i>
<i>Die Grundlagen der Korrelationen</i>	<i>36</i>
<i>Anwendung von Korrelationen in Risiko Simulator</i>	<i>37</i>
<i>Die Effekte von Korrelationen in Monte-Carlo-Simulationen.....</i>	<i>38</i>
<i>Präzisions- und Fehlerkontrolle.....</i>	<i>40</i>
<i>Die Vorausberechnungsstatistiken begreifen.....</i>	<i>43</i>
<i>Vermessen des Verteilungszentrums — das erste Moment.....</i>	<i>43</i>
<i>Vermessen der Verteilungsdispersion — das zweite Moment.....</i>	<i>43</i>
<i>Vermessen der Verteilungsschiefe — das dritte Moment.....</i>	<i>45</i>
<i>Vermessen der katastrophalen Schwanzereignissen in einer Verteilung — das vierte Moment</i>	<i>46</i>
<i>Wahrscheinlichkeitsverteilungen für Monte-Carlo-Simulationen begreifen</i>	<i>48</i>

<i>Diskrete Verteilungen</i>	51
<i>Bernoulli oder Ja/Nein Verteilung</i>	51
<i>Binomialverteilung</i>	52
<i>Diskrete Uniform</i>	53
<i>Geometrische Verteilung</i>	53
<i>Hypergeometrische Verteilung</i>	54
<i>Negative Binomialverteilung</i>	55
<i>Pascal-Verteilung</i>	56
<i>Poisson-Verteilung</i>	57
<i>Kontinuierliche Verteilungen</i>	58
<i>Arcussinus-Verteilung</i>	58
<i>Betaverteilung</i>	58
<i>Multiplikative Beta versetzte Verteilung</i>	59
<i>Cauchy- oder Lorentz- oder Breit-Wigner-Verteilung</i>	60
<i>Chi-Quadrat Verteilung</i>	60
<i>Doppelte Log-Verteilung</i>	61
<i>Dreiecksverteilung</i>	62
<i>Erlang-Verteilung</i>	63
<i>Exponentialverteilung</i>	63
<i>Exponentialverteilung Versetzte</i>	64
<i>Extremwert- oder Gumbel-Verteilung</i>	64
<i>F-Verteilung oder Fisher-Snedecor-Verteilung</i>	65
<i>Gammaverteilung (Erlang-Verteilung)</i>	66
<i>Kosinus-Verteilung</i>	67
<i>Laplace-Verteilung</i>	68
<i>Logistische Verteilung</i>	68
<i>Lognormalverteilung</i>	69
<i>Lognormale Versetzte Verteilung</i>	70
<i>Normalverteilung</i>	70
<i>Parabolische-Verteilung</i>	71
<i>Pareto-Verteilung</i>	72
<i>Pearson-V-Verteilung</i>	73
<i>Pearson-VI-Verteilung</i>	73
<i>PERT-Verteilung</i>	74
<i>Potenzverteilung</i>	75
<i>Multiplikative Potenzverteilung versetzte</i>	76
<i>Studentsche t Verteilung</i>	76

Uniformverteilung	77
Weibull-Verteilung (Rayleigh-Verteilung).....	77
Multiplikative Weibull- und Rayleighverteilungen versetzte	78
3. VORAUSBERECHNUNG.....	80
Verschiedenen Typen von Vorausberechnungsverfahren	80
Das Vorausberechnungs-Tool in Risiko Simulator ausführen.....	85
Zeitreihenanalyse.....	86
Multivariate Regression.....	90
Stochastische Vorausberechnung	94
Nichtlineare Extrapolation	98
Box-Jenkins ARIMA Fortgeschrittene Zeitreihen	100
AUTO ARIMA (Box-Jenkins ARIMA Fortgeschrittene Zeitreihen)	106
Grund-Ökonometrie.....	107
J-S Kurven Vorausberechnungen	109
GARCH Volatilitätsvorausberechnungen.....	111
Markov-Ketten	113
Maximale Wahrscheinlichkeitsmodelle (MLE) auf Logit, Probit und Tobit.....	114
Spline (Kubischer Spline Interpolation und Extrapolation)	117
4. OPTIMIERUNG.....	120
Optimierungsmethodologien.....	120
Optimierung mit kontinuierlichen Entscheidungsvariablen	123
Optimierung mit diskreten ganzzahligen Variablen	128
Effiziente Grenze und fortgeschrittene Optimierungseinstellungen	133
Stochastische Optimierung	134
5. ANALYTISCHE TOOLS IN RISIKO SIMULATOR.....	141
Tornado und Sensibilität Tools in der Simulation	141
Sensibilitätsanalyse.....	150
Verteilungsanpassung: Einzel-Variable und Mehrfach-Variablen.....	153
Bootstrap-Simulation	158
Hypothesentest	161
Daten extrahieren und Simulationsergebnisse speichern	163
Ein Bericht erstellen	164

<i>Regressions- und Vorausberechnungs-Diagnosetool</i>	166
<i>Statistische Analyse Tool</i>	174
<i>Verteilungsanalyse Tool</i>	180
<i>Szenarioanalyse Tool</i>	183
<i>Segmentierung Clustering Tool</i>	185
RISIKO SIMULATOR 2011/2012 NEUE WERKZEUGE	186
<i>Zufallszahl Generierung, Monte Carlo versus Latin Hypercube, und Korrelation</i>	
<i>Kopula Methoden</i>	186
<i>Daten Saisonalbereinigung und Trendbereinigung</i>	187
<i>Hauptkomponentenanalyse</i>	189
<i>Strukturbruchanalyse</i>	190
<i>Trendlinie Voraussagen</i>	191
<i>Modellprüfungstool</i>	192
<i>Perzentiles Verteilungsanpassungs-Tool</i>	193
<i>Verteilungscharts und Tabellen: Wahrscheinlichkeitsverteilungen Tool</i>	195
<i>ROV BizStats</i>	199
<i>Neuronales Netzwerk und Kombinatorische Fuzzy Logic Voraussagen Methodologien</i>	204
<i>Optimierer Goal Seek</i>	208
<i>Einzel Variable Optimierer</i>	208
<i>Genetische Algorithmus Optimierung</i>	209
<i>Modul ROV Entscheidungsbaum</i>	211
<i>Simulationsmodellierung</i>	213
<i>Bayessche Analyse</i>	213
<i>Erwartungswert der perfekten Information, Minimax und Maximin Analyse,</i>	
<i>Risikoprofile und Wert der unvollständigen Information</i>	214
<i>Sensibilität</i>	215
<i>Szenarietabellen</i>	215
<i>Generation der nutzenfunktion</i>	215
HILFBREICHE TIPPS UND TECHNIKEN	224
<i>TIPPS: Annahmen (Festgesetzte Input- Annahme Benutzeroberfläche)</i>	224
<i>TIPPS: Kopieren und Einfügen TIPS: Copy and Paste</i>	224
<i>TIPPS: Korrelationen</i>	225
<i>TIPPS: Datendiagnostik und Statistische Analyse</i>	225

<i>TIPPS: Verteilungsanalyse, Charts und Wahrscheinlichkeitstabellen</i>	<i>226</i>
<i>TIPPS: Effiziente Grenze.....</i>	<i>226</i>
<i>TIPPS: Vorassagezellen</i>	<i>226</i>
<i>TIPPS: Voraussageschart</i>	<i>227</i>
<i>TIPPS: Voraussagen</i>	<i>227</i>
<i>TIPPS: Voraussagen: ARIMA</i>	<i>227</i>
<i>TIPPS: Vorassagen: Basic Ökonometrie.....</i>	<i>227</i>
<i>TIPPS: Vorassagen: Logit, Probit und Tobit</i>	<i>228</i>
<i>TIPPS: Voraussagen: Stochastische Prozesse</i>	<i>228</i>
<i>TIPPS: Voraussagen: Trendlinien.....</i>	<i>228</i>
<i>TIPPS: Funktionsaufrufe.....</i>	<i>228</i>
<i>TIPPS: Einstiegsübungen und Einstiegsvideos</i>	<i>229</i>
<i>TIPPS: Hardware ID.....</i>	<i>229</i>
<i>TIPPS: Latin Hypercube Sampling (LHS) gegen Monte Carlo Simulation (MCS)</i>	<i>229</i>
<i>TIPPS: Online Ressourcen TIPS: Online Resources.....</i>	<i>230</i>

1. EINLEITUNG

Willkommen zur Software RISIKO SIMULATOR

Der **Risiko Simulator** ist eine Software für Monte-Carlo-Simulation, Vorausberechnung (Prognosen) und Optimierung. Die Software ist in Microsoft .NET C# geschrieben und funktioniert als Add-in zusammen mit Excel. Diese Software ist auch kompatibel und wird oft verwendet mit der Software Real-Optionen Super-Verband-Löser (SLS) und der Software Belegschaftsaktienoptionen Bewertungs-Toolkit (ESOV), auch von Real Options Valuation, Inc. entwickelt. Bitte bemerken Sie, dass trotz unserer Bemühung sehr gründlich in diesem Handbuch vorzugehen, dieses Handbuch ist auf keinen Fall ein Ersatz für die Trainings-DVD, für Live-Trainingkurse und für Bücher verfasst vom Erschaffer dieser Software (z.B., *Real Options Analysis*, 2nd Edition, Wiley Finance 2005; *Modeling Risk: Applying Monte Carlo Simulation*, *Real Options Analysis, Forecasting, and Optimization*, 2nd Edition, Wiley 2010; und *Valuing Employee Stock Options (2004 FAS 123R)*, Wiley Finance, 2004, alle von Dr. Johnathan Mun). Bitte besuchen Sie unsere Webseite unter www.realoptionsvaluation.com für mehr Informationen bezüglich dieser Einzelheiten.

Die Software **Risiko Simulator** beinhaltet die folgenden Module:

- Monte-Carlo-Simulation (führt parametrische und nicht-parametrische Simulationen von 42 Wahrscheinlichkeitsverteilungen mit verschiedenen Simulationsprofile, gestutzten und korrelierten Simulationen, angepassten Verteilungen, Präzisions- und Fehlerkontrollierten Simulationen und vielen anderen Algorithmen aus)
- Vorausberechnung (führt Box-Jenkins ARIMA, Mehrfachregressionen, nicht-lineare Extrapolationen, stochastische Prozesse und Zeitreihenanalysen aus)
- Optimierung unter Ungewissheit (führt Optimierungen unter Verwendung von diskreten ganzzahligen und kontinuierlichen Variablen für Portfolio- und Projektoptimierung mit und ohne Simulation aus)
- Modellierung und analytische Tools (führt Tornado-, Spinnennetz- und Sensibilitätsanalysen, ebenso wie Bootstrap-Simulationen, Hypothesentesten, Verteilungsanpassungen, usw., aus)

Die Software **Real-Optionen SLS** wird verwendet, um einfache und komplexe Optionen zu berechnen und schließt die Fähigkeit angepasste Optionsmodelle zu kreieren ein. Die Software **Real-Optionen SLS** beinhaltet die folgenden Module:

- Einzel-Aktivum SLS (zur Lösung von Abbruchs-, Chooser-, Kontraktions-, Verschiebungs- und Expansionsoptionen, sowie auch zur Lösung von angepassten Optionen)
- Mehrfach-Aktiva und Mehrphasen SLS (zur Lösung von mehrphasig sequenziellen Optionen, Optionen mit mehrfachen unterliegenden Aktiva und Phasen, Kombinationen von mehrphasig sequenziellen Abbruchs-, Chooser-, Kontraktions-, Verschiebungs-,

Expansions- und Switching-Optionen; auch zur Lösung von angepassten Optionen verwendbar)

- Multinomial SLS (zur Lösung von trinomialen Rückkehr-zum-Mittelwert-Optionen, quadrinomialen Sprung-Diffusion-Optionen und pentanomialen Regebogen-Optionen)
- Excel Add-in Funktionen (zur Lösung aller obigen Optionen plus geschlossenen Form Modellen und angepassten Optionen in einer auf Excel basierenden Umgebung)

Installationsvoraussetzungen und -prozeduren

Um die Software zu installieren, folgen Sie den auf dem Bildschirm erscheinenden Anweisungen. Die Mindestvoraussetzungen für diese Software sind:

- Pentium IV Prozessor oder spätere Version (Dual-Core empfohlen).
- Windows XP, Vista oder Windows 7.
- Microsoft Excel XP, 2003, 2007, 2010 oder spätere Version.
- Microsoft .NET Framework 2.0 oder 3.0 oder spätere Version.
- 500MB verfügbarer Speicher.
- 2GB RAM Minimum (2-4GB empfohlen).
- Administratorrechte zur Softwareinstallation.

Die meisten neuen Rechner sind mit vorinstalliertem Microsoft .NET Framework 2.0/3.0 erhältlich. Wenn allerdings eine Fehlermeldung bezüglich der Anforderung von .NET Framework während der Installation von Risiko Simulator erscheint, beenden Sie die Installation. Dann installieren Sie die entsprechende Software .NET Framework, welche in der CD beigelegt ist (wählen Sie Ihre eigene Sprache). Schließen Sie die Installation von .NET ab, starten Sie den Computer neu und installieren Sie erneut die Software Risiko Simulator. Eine Standard 10-Tage Probelizenzdatei ist der Software beigelegt. Um eine vollständige Firmenlizenz zu erhalten, kontaktieren Sie bitte Real Options Valuation, Inc. unter admin@realoptionsvaluation.com oder rufen Sie (925) 271-4438 an oder besuchen Sie unsere Webseite unter www.realoptionsvaluation.com. Besuchen Sie bitte unsere Webseite und klicken Sie auf DOWNLOAD, um die neuste Softwareversion zu bekommen, oder klicken Sie auf die Verknüpfung FAQ, um aktualisierte Informationen über die Lizenzvergabe oder Installationsthemen und -verbesserungen zu erhalten.

Lizenzvergabe

Nach der Installation der Software und dem Erwerb einer vollständigen zur Verwendung der Software erforderlichen Lizenz, müssen Sie uns Ihre Hardware-ID per E-Mail senden, sodass wir eine Lizenzdatei für Sie generieren können. Befolgen Sie die nachstehenden Anweisungen:

Für **Windows XP** mit Excel XP, Excel 2003 oder Excel 2007:

- Als erstes, in Excel klicken Sie auf **Risiko Simulator | Lizenz**, kopieren Sie Ihre elfstellige alphanumerische HARDWARE-ID und senden Sie uns diese per E-Mail an die folgende Adresse: admin@realoptionsvaluation.com. (Sie können auch die Hardware-ID auswählen und Sie durch einen Rechtsklick kopieren oder auf die Verknüpfung Hardware-ID mailen klicken). Sobald wir diese ID erhalten haben, werden wir Ihnen eine neu generierte permanente Lizenz per E-Mail senden. Wenn Sie diese Lizenzdatei erhalten haben, speichern Sie diese einfach auf Ihre Festplatte. Starten Sie Excel und klicken Sie auf **Risiko Simulator | Lizenz**. Dann klicken Sie auf **Lizenz installieren** und zeigen Sie auf diese neue Lizenzdatei. Starten Sie Excel neu und der Vorgang ist beendet. Das gesamte Verfahren dauert weniger als eine Minute und aktiviert die vollständige Lizenz zur Ausführung der Software.

Für **Windows Vista, Windows 7** mit Excel XP, Excel 2003 oder Excel 2007:

- Als erstes, starten Sie Excel 2007 innerhalb Windows Vista, gehen Sie zur Menüleiste Risiko Simulator und klicken Sie auf der Ikone **Lizenz** oder auf **Risiko Simulator | Lizenz**. Jetzt kopieren und senden Sie per E-Mail Ihre elfstellige alphanumerische HARDWARE-ID (Sie können auch die Hardware-ID auswählen und Sie durch einen Rechtsklick kopieren oder auf die Verknüpfung „Hardware-ID mailen“ klicken) an die folgende Adresse: admin@realoptionsvaluation.com. Sobald wir diese ID erhalten haben, werden wir Ihnen eine neu generierte permanente Lizenz per E-Mail senden. Wenn Sie diese Lizenzdatei erhalten haben, speichern Sie diese einfach auf Ihre Festplatte. Starten Sie Excel und klicken Sie auf **Risiko Simulator | Lizenz** oder auf die Ikone **Lizenz**. Dann klicken Sie auf **Lizenz installieren** und zeigen Sie auf diese neue Lizenzdatei. Starten Sie Excel neu und der Vorgang ist beendet. Das gesamte Verfahren dauert weniger als eine Minute und aktiviert die vollständige Lizenz zur Ausführung der Software.

Sobald die Installation beendet ist, starten Sie Microsoft Excel. Wenn die Installation erfolgreich war, müssten sowohl ein zusätzliches Element “*Risiko Simulator*” in der Menüleiste in Excel XP/2003 oder unter der ADD-IN Gruppe in Excel 2007 als auch eine neue Ikoneleiste in Excel erscheinen, wie im Bild 1.1 angezeigt. Außerdem erscheint ein Begrüßungsbildschirm, siehe Bild 1.2, welcher anzeigt, dass die Software funktioniert und in Excel geladen ist. Das Bild 1.3 zeigt auch die Symbolleiste von Risk Simulator. Wenn diese Elemente in Excel vorhanden sind, können sie beginnen die Software zu verwenden. Die folgenden Abschnitte geben Ihnen schrittweise Anweisungen zur Verwendung der Software.

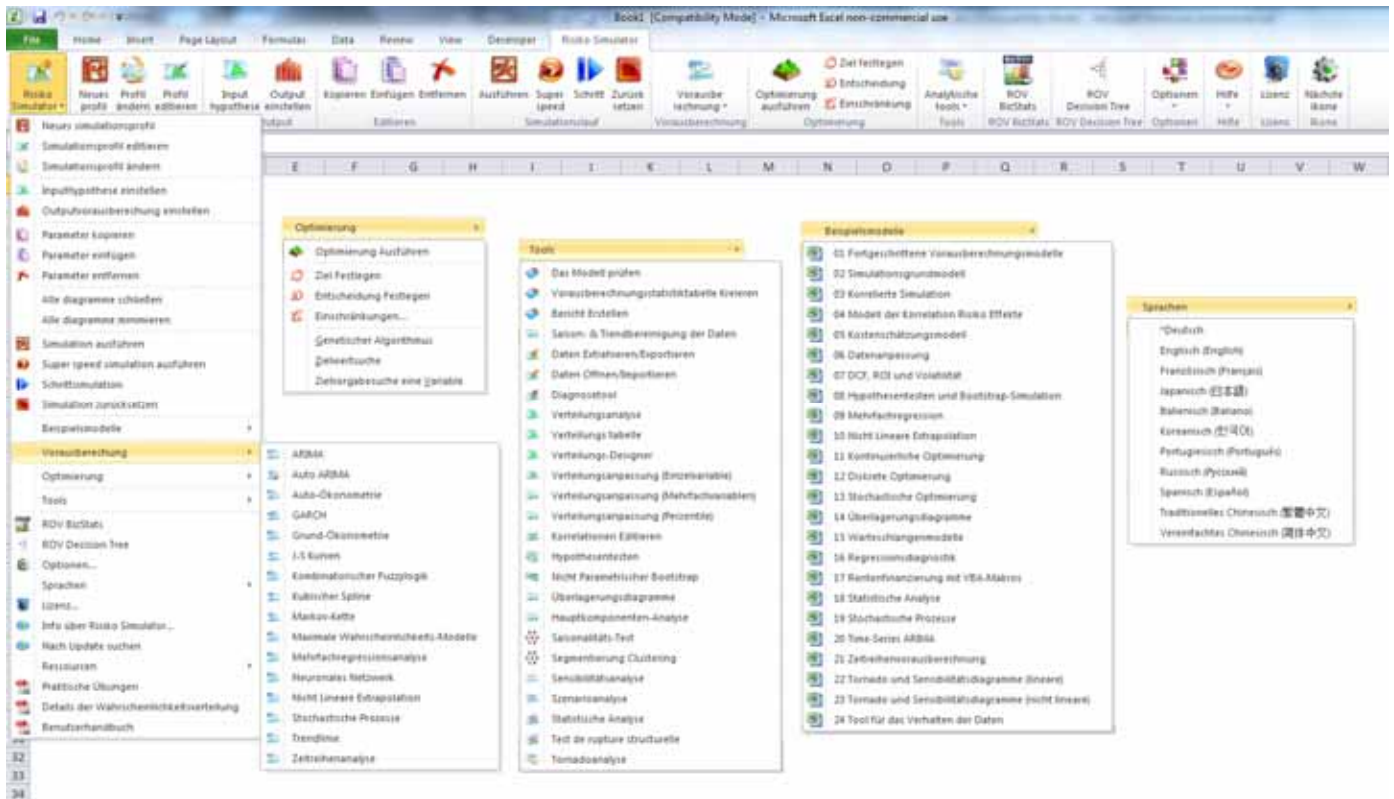


Bild 1.1—Risiko Simulator Menü und Symbolleiste in Excel 2007



Bild 1.2—Risiko Simulator Begrüßungsbildschirm

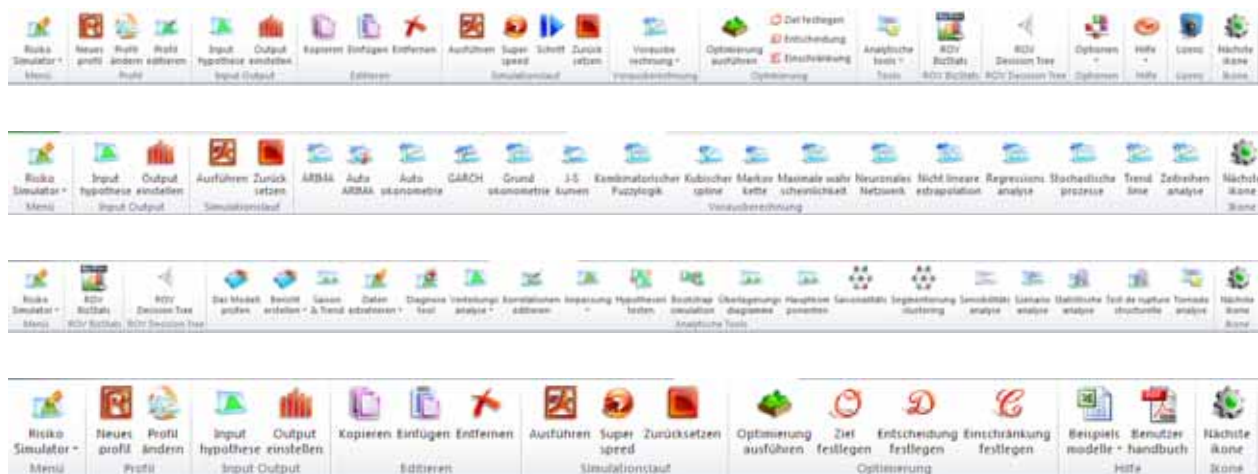


Bild 1.3—Risiko Simulator Ikonen-Symbolleiste in 2007/2012

WAS GIBT'S NEUES IN VERSION 2011/2012

Ein detailliertes Leistungsverzeichnis des Risiko Simulators

In der folgenden Aufzählung sind die wichtigsten Einsatzmöglichkeiten des Risiko Simulators, wobei die neuesten Ergänzungen der Version 2011/2012 hervorgehoben sind.

Allgemeine Einsatzmöglichkeiten

1. In 11 Sprachen verfügbar—Englisch, Französisch, Deutsch, **Italienisch**, Japanisch, **Koreanisch**, Portugiesisch, Spanish, **Russisch**, Chinesisch (Kurzzeichen) und **Chinesisch (Langzeichen)**.
2. *ROV Entscheidungsbaum wird verwendet, um Entscheidungsbaummodelle zu erstellen und bewerten. Zusätzliche fortgeschrittene Methodologien und Analytiken sind ebenfalls enthalten:*
 - **Entscheidungsbaummodelle**
 - **Monte-Carlo-Risikosimulation**
 - **Sensibilitätsanalyse**
 - **Szenarienanalyse**
 - **Bayessche-Analyse (gesamte und a posteriori Wahrscheinlichkeitsaktualisierung)**
 - **Erwartungswert der Information**
 - **MINIMAX**
 - **MAXIMIN**
 - **Risikoprofile**
3. Bücher—analytische Theorie, Anwendung und Fallbeispiele werden von 10 Büchern unterstützt.
4. Kommentierte Zellen— die kommentierte Zellen an- oder ausschalten, und entscheiden ob Sie Zellkommentare an allen Eingabe-Annahmen, Ausgabe-Prognosen und Entscheidungs-Variablen zeigen möchten.
5. Detaillierte Beispielmmodelle— 24 Beispielmmodelle im Risiko Simulator und mehr als 300 Modelle im Modellierung Toolkit (Modeling Toolkit).
6. Detaillierte Berichte—Alle Analysen verfügen über detaillierten Berichten.
7. Detailliertes Anwenderhandbuch—Ein Schritt-für-Schritt Anwenderhandbuch.
8. Flexible Lizenzierung—In der Lage bestimmte Funktionen an-oder auszuschalten um Ihre Risikoanalyse Erfahrungswerte anzupassen (einzurichten). Zum Beispiel, wenn Sie sich nur für die Prognose Tools im Risiko Simulator interessieren, können Sie vielleicht eine Sonder-Lizenz bekommen, welche nur die Prognose Tools aktiviert während die andere Module deaktiviert sind, und somit einige Software Kosten sparen.
9. Flexible Voraussetzungen—Funktioniert mit Windows 7, Vista und XP; integriert sich mit Excel 2010, 2007, 2003; und läuft auch unter MAC Betriebssysteme mit virtuellen Maschinen.
10. Vollständig anpassbare Farben und Diagramme—neigen, 3D, Farbe, Chart-Typ, und vieles mehr!

11. Praktische Anwendungen—Detaillierte Schritt-für-Schritt Anleitung für den Betrieb des Risiko Simulators, inklusive Hilfe in der Auswertung der Ergebnisse.
12. Multiples Ausschneiden und Einfügen—lässt Hypothesen, Annahmen Entscheidungs-Variablen und Prognosen kopieren und einfügen zu.
13. Das Profiling—lässt die Erzeugung mehrere Profile in einem einzelnen Model zu. (Verschiedene Szenarien der Simulationsmodellen können erstellt, vervielfacht, editiert und im Singelmodus ausgeführt werden).
14. Revidierte Symbole in Excel 2007/2010— eine komplett neu definierte Symbolleiste, die noch intuitive und benutzerfreundliche ist. Es stehen vier Piktogramm-Sätze zur Verfügung, die sich den meisten Bildschirmauflösungen anpassen (1280 x 760 und höher).
15. Rechtsklick Verknüpfungen—alle Risiko Simulator Tools und Menüs mit einem rechten Mausklick abrufen.
16. ROV Software Integration—funktioniert gut mit anderen ROV Software inklusive Real Options SLS, Modeling Toolkit, Basel Toolkit, ROV Compiler, ROV Extractor und Evaluator, ROV Modeler, ROV Valuator, ROV Optimizer, ROV Dashboard, ESO Valuation Toolkit, und andere!
17. RS Funktionen in Excel—RS Funktionen einfügen um Annahmen und Prognosen einzurichten, und Rechtsklick Unterstützung im Excel.
18. Troubleshooter—Dieses Tool ermöglicht Ihnen Ihre Software zu reaktivieren, Ihre System Anforderungen zu überprüfen, die Hardware ID zu bekommen, und andere.
19. Turbo Speed Analyse—Diese neue Fähigkeit lässt Prognosen und andere Analyse Tools in auffallend schnellen Geschwindigkeiten laufen (in Version 5.2 verbessert). Die Analysen und Ergebnisse bleiben gleich, werden aber jetzt sehr schnell errechnet und in Berichten generiert.
20. Web Ressourcen, Fallstudien und Videos—kostenfreie Modelle herunterladen, Einstiegs-Videos, Fallbeispiele, White Papers, und andere Materialien von unserer Website.

Simulationsmodell

21. 6 Zufallsgeneratoren—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 Abfragemethoden—Monte Carlo und Latin Hypercube.
23. 3 Korrelations-Kopula—Anwendung von Normal Copula, T Copula, und Quasi-Normal Copula für korrelierte Simulationen.
24. 42 Wahrscheinlichkeitsverteilungen—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. Perzentilwerte anwenden als Alternative Art um Parameter einzugeben.

26. Vertriebsschema ohne vorgegebene Parameter—Ihren eigenen Vertrieb bestimmen, historische Simulationen durchführen, und die Delphi Methode anwenden
27. Vertrieb Trunkierung—Daten Grenzen ermöglichen.
28. Excel Funktionen—Annahmen und Prognosen einstellen unter Verwendung der Funktionen innerhalb Excel.
29. Multidimensionale Simulation—Simulation der unbestimmten Eingabeparameter.
30. Präzisionskontrolle—stellt fest, ob die Anzahl der durchgeführten Simulation Testläufe ausreichend ist.
31. Super Geschwindigkeits-Simulation—führt 100.000 Testläufe in wenigen Sekunden aus.

Prognose Modul

32. ARIMA—autoregressive, integrierte, gleitender Durchschnittsmodelle ARIMA (P,D,Q).
33. Auto ARIMA—führt die häufigsten Kombinationen der ARIMA aus, um das passendste Model zu finden.
34. Auto Ökonometrie—führt tausende Model-Kombinationen und Permutationstests aus, um das passendste Model for die bestehenden Daten zu erzielen (linear, nicht linear, wechselwirkend, Schwankungen, Quote/Preis, Differenz).
35. Basic Ökonometrie—ökonometrisch und linear/nicht linear und wechselwirkenden Regressionsmodelle.
36. Kubistische Spline—die nicht lineare Interpolation und Extrapolation.
37. GARCH—Volatilitätsberechnungen mittels verallgemeinerten, autoregressiven, konditionalen Heteroskedastie Modellen: *GARCH, GARCH-M, TGARCH, TGARCH-M, EGARCH, EGARCH-T, GJR-GARCH, und GJR-TGARCH.*
38. J-Kurve—exponentielle J Kurven.
39. Limitierte abhängige Variablen—*Logit, Probit und Tobit.*
40. Markov Chains—zwei konkurrierende Elemente über Zeit- und Marktanteil Prognosen.
41. Mehrfachregression—reguläre linear und nicht linear Regression mit schrittweise Methodologien (vor, zurück, Korrelation, vor-zurück).
42. Nicht lineare Extrapolation—nicht-lineare Zeitreihenprognose.
43. S Kuurve—logistische S Kurven.
44. Zeitreihenanalyse—8 Zeitreihen Zerlegungsmodelle für die Prognose von Ebenen, Tendenzen und Saisonbewegungen.
45. *Trendlinien—Prognose und Installation mittels linear, nicht linear polynomisch, Leistung, logarithmisch, exponentiell, und geglätteter Mittelwert mit guter Anpassung.*

Optimierungsmodell

46. Linear Optimierung—mehrphasig Optimierung und allgemeine linear Optimierung.
47. Nicht linear Optimierung—detaillierte Ergebnisse inklusive Hessematrizen, LaGrange Funktionen und mehr.
48. Statische Optimierung—Schnelldurchläufe für kontinuierliche, ganzzahlig und binäre Optimierungen.

49. Dynamische Optimierungen—Simulation mit Optimierung.
50. Stochastische Optimierung—quadratische, tangentiale, zentrale, vorwärts, Konvergenzkriterien.
51. Effiziente Grenze—Kombinationen von stochastischen und dynamischen Optimierungen an multivariaten effizienten Grenzen.
52. **Genetische Algorithmen**—angewendet für eine Vielzahl der Probleme bei der Optimierung.
53. Mehrphasige Optimierung—lokale versus globale Optimum testen, um bessere Kontrolle über die Optimierungsabläufe zu ermöglichen, und die Genauigkeit und Abhängigkeit der Ergebnisse zu verbessern.
54. Perzentile und konditionale Mittelwerte—zusätzliche Statistiken für stochastische Optimierung, inklusive Perzentile sowie konditionale Mittelwerte, die für die Berechnung des konditionalen Wertes bei Risiko Maßnahmen entscheidend sind
55. **Suchalgorithmus**—einfache, schnelle und effiziente Suchalgorithmen für grundlegenden einzelnen Entscheidungsvariablen und Zielfindungsanwendungen.
56. Super Geschwindigkeit Simulation in dynamischen und stochastischen Optimierung—lässt die Simulation mit super Geschwindigkeit laufen während der Integration mit der Optimierung.

Analytisches Hilfsprogramm Modul

57. **Prüfmodul**—testet die häufigsten Fehler in Ihrem Model aus.
58. Korrelations-Editor—ermöglicht die direkte Eingabe und Editierung großer Korrelationsmatrizen.
59. Berichterstellung—Automatisierung der Berichterstellung von Annahmen und Prognosen in einem Model.
60. Erstellung statistischer Berichte—erstellt komparativen Bericht aller Prognosen Statistiken.
61. Daten Diagnose—führt Testläufe aus zu Heteroskedastie, Mikronumerosität, Ausreiser, Nichtlinearität, Autokorrelation, Normalität, Sphärizität, Nicht-Stationarität, Multikollinearität und Korrelationen.
62. Daten Extraktion und Export—Daten entpacken für Excel oder flache Textdateien und Risk Sim Dateien, führt statistische Berichte und Prognose Ergebnisberichte durch.
63. Daten öffnen und importieren—ruft vorige Simulationslauergebnisse ab.
64. **Saisonbereinigung und Trend-Entfernung**—saisonbereinigt und entfernt Trends von Ihren Daten.
65. Verteilungsanalyse—berechnet exakt PDF, CDF und ICDF von allen 42 Verteilungen und erstellt Wahrscheinlichkeitstabellen.
66. Verteilungs-Designer—Ihre eigenen maßgeschneiderten Verteilungen erstellen.
67. Verteilungs-Anpassung (multipel) - führt multipel Variablen simultan aus, weist Korrelation und Korrelation Signifikanz aus.
68. Verteilungs-Anpassung (einzeln)—Kolmogorov-Smirnov und Chi-Quadrat Prüfungen der fortlaufenden Verteilungen, komplett mit Berichten und Verteilungsannahmen.
69. Hypothesenprüfung—prüfen, ob zwei Prognosen statistisch ähnlich oder verschieden sind.

70. Nichtparametrisches Bootstrapping—Simulation der Statistiken um Präzision und Akkurateesse der Ergebnissen zu erreichen.
71. Overlay Charts/Folien—vollständig anpassbaren Overlay Charts/Folien von Annahmen und Prognosen zugleich (CDF, PDF, 2D/3D Chart/Folie Varianten).
72. **Hauptkomponentenanalyse**—prüft die besten Prädiktor-Variablen und Methoden um das Datenfeld zu reduzieren.
73. Szenarioanalyse—hunderte und tausende statischer zweidimensionalen Szenarien.
74. Saisonalitäts-Test—Tests für verschiedene saisonbedingten Verzögerungen.
75. Segmentations-Clustering—teilt die Daten in statistischen Gruppen um Ihre Daten zu segmentieren.
76. Sensitivitätsanalyse—dynamische Sensitivität (Simultananalyse).
77. **Strukturelles Abbruch/Unterbrechungstest**—prüft ob Ihre Zeitreihen Daten statistisch strukturierte Abbrüche/Unterbrechungen.
78. Tornado-Analyse—statische Perturbation der Empfindlichkeiten, Spider und Tornado Analyse, und Szenario Tabellen.

Statistiken und BizStats Modul Statistics

79. **Perzentiles Verteilungs-Anpassung**—Perzentile und Optimierung verwenden um die passendste Verteilung zu finden.
80. **Wahrscheinlichkeitsverteilung—Charts und Tabellen**—lässt 45 Wahrscheinlichkeitsverteilungen, ihre vier Momente CDF, ICDF, PDF, Charts, Overlay Multipel-Verteilungs-Charts laufen, und erstellt Wahrscheinlichkeitsverteilung Tabellen.
81. Statistische Analyse - erklärende Statistiken, Verteilungs-Anpassung, Histogramme, Charts, nichtlineare Extrapolation, Normalitätstest, stochastische Parameter-Bewertung, Zeitreihen Prognosen, Trendlinie Projektionen, usw.
82. **ROV BIZSTATS**—mehr als 130 Wirtschaftsstatistik und analytische Modelle:

Absolute Values, ANOVA: Randomized Blocks Multiple Treatments, ANOVA: Single Factor Multiple Treatments, ANOVA: Two Way Analysis, ARIMA, Auto ARIMA, Autocorrelation and Partial Autocorrelation, Autoeconometrics (Detailed), Autoeconometrics (Quick), Average, Combinatorial Fuzzy Logic Forecasting, Control Chart: C, Control Chart: NP, Control Chart: P, Control Chart: R, Control Chart: U, Control Chart: X, Control Chart: XMR, Correlation, Correlation (Linear, Nonlinear), Count, Covariance, Cubic Spline, Custom Econometric Model, Data Descriptive Statistics, Deseasonalize, Difference, Distributional Fitting, Exponential J Curve, GARCH, Heteroskedasticity, Lag, Lead, Limited Dependent Variables (Logit), Limited Dependent Variables (Probit), Limited Dependent Variables (Tobit), Linear Interpolation, Linear Regression, LN, Log, Logistic S Curve, Markov Chain, Max, Median, Min, Mode, Neural Network Forecasts, 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. MONTE-CARLO-SIMULATION

Die Monte-Carlo-Simulation, genannt nach der Spielcasinostadt von Monaco, ist eine sehr leistungsfähige Methodologie. Die Simulation öffnet dem Fachmann die Tür zur sehr einfachen Lösung von schwierigen und komplexen aber praktischen Problemen. Monte-Carlo kreiert artifizielle Zukünfte, indem sie Tausende und sogar Millionen von Probepfaden von Ereignissen generiert und ihre vorherrschenden Eigenschaften examiniert. Für Analysten in einer Firma ist die Teilnahme an fortgeschrittenen postakademischen Mathekursen nicht nur logisch oder praktisch. Ein brillanter Analyst würde alle Ihm/Ihr zur Verfügung stehenden Instrumente benutzen, um die gleiche Antwort auf die möglichst einfache und praktischste Art zu bekommen. In allen Fällen, wenn korrekt modelliert, liefert die Monte-Carlo-Simulation außerdem ähnliche Antworten wie mathematisch elegantere Methoden. Also, was ist die Monte-Carlo-Simulation und wie funktioniert sie?

Was ist die Monte-Carlo-Simulation?

Die Monte-Carlo-Simulation in ihrer einfachsten Form ist ein Zufallszahlengenerator, der für die Vorausberechnung, Schätzung und Risikoanalyse nützlich ist. Eine Simulation berechnet zahlreiche Szenarien eines Modells, indem sie wiederholt Werte aus einer benutzervordefinierten *Wahrscheinlichkeitsverteilung* für die ungewissen Variablen auswählt und diese Werte für das Modell verwendet. Alle diese Szenarien produzieren dazugehörige Ergebnisse, wobei jedes Szenario eine *Vorausberechnung* (Prognose) haben kann. Vorausberechnungen sind Ereignisse (normalerweise mit Formeln oder Funktionen), die Sie als wichtige Outputs des Modells definieren. Dies sind normalerweise Ereignisse wie Endsummen, Nettogewinn oder Bruttoausgaben.

Zur Vereinfachung, denken Sie an die Methode der Monte-Carlo-Simulation wie das wiederholte Herausnehmen mit Zurücklegung von Golfbällen aus einem großen Korb. Die Größe und Form des Korbs hängt von der *Inpu*thypothese der Verteilung ab (z.B., eine Normalverteilung mit einem Mittelwert von 100 und einer Standardabweichung von 10, gegen eine Uniformverteilung oder eine Dreiecksverteilung), wobei einige Körbe tiefer oder symmetrischer sind als andere, was bedeutet, dass einige Bälle häufiger als andere herausgenommen werden. Die Anzahl der wiederholt herausgenommenen Bälle hängt von der Anzahl der simulierten *Probeversuche* ab. Für ein großes Modell mit mehrfachen verwandten Hypothesen, stellen Sie sich das große Modell wie einen sehr großen Korb vor, in dem viele Minikörbe stecken. Jeder Minikorb hat seine eigene Menge von herum springenden Golfbällen. Gelegentlich stehen diese Minikörbe in Verbindung zueinander (wenn es eine *Korrelation* zwischen den Variablen gibt) und die Golfbälle springen im Zusammenhang miteinander herum, während andere unabhängig voneinander herumspringen. Die Bälle, die jedes Mal aus diesen Interaktionen innerhalb des Modells (der große zentrale Korb) herausgenommen werden, werden tabelliert und aufgezeichnet, was ein *Vorausberechnung*soutputergebnis der Simulation liefert.

Ein Überblick der Software

Die Software *Risiko Simulator* hat mehrere Anwendungen, einschließlich die Monte-Carlo-Simulation, die Vorausberechnung (Prognose), die Optimierung und die Risikoanalytik.

- ② Das *Modul Simulation* erlaubt Ihnen Folgendes: Simulationen in Ihren existierenden Excelbasierenden Modellen auszuführen; Simulationsvorausberechnungen (Verteilungen von Ergebnissen) zu generieren und extrahieren; Verteilungsanpassungen durchzuführen (automatisch die bestpassende statistische Verteilung zu finden); Korrelationen (Beibehaltung von Verhältnissen zwischen simulierten Zufallsvariablen) zu berechnen; Empfindlichkeiten zu identifizieren (Tornado- und Sensibilitätsdiagramme zu erstellen); statistische Hypothesen zu testen (statistische Unterschiede zwischen Vorausberechnungspaaren zu finden); Bootstrapsimulation auszuführen (die Robustheit der Ergebnisstatistiken zu testen); angepasste und nicht-parametrische Simulationen auszuführen (Simulationen unter Verwendung von historischen Daten ohne irgendwelche Verteilungen oder deren Parameter zu spezifizieren, um Vorausberechnungen ohne Daten auszuführen oder um Expertenmeinungsprognosen anzuwenden).
- ② Das *Modul Vorausberechnung* kann verwendet werden, um Folgendes zu generieren: automatische Zeitreihenvorausberechnungen (mit und ohne Saisonalität und Trend); multivariate Regressionen (Modellierung von Verhältnissen zwischen Variablen); nichtlineare Extrapolationen (Kurvenanpassung); stochastische Prozesse (Irrfahrt, Rückkehr zum Mittelwert, Sprung-Diffusion und kombinierte Prozesse); Box-Jenkins ARIMA (ökonometrische Vorausberechnungen); Auto ARIMA; Grund-Ökonometrie und Auto- Ökonometrie (Modellierung von Verhältnissen und Erstellung von Vorausberechnungen); exponentielle J-Kurven; logistische S-Kurven; GARCH Modelle und deren mehrfache Variationen (Modellierung und Vorausberechnung der Volatilität); maximale Wahrscheinlichkeitsmodelle für limitierte abhängige Variablen (Logit, Tobit und Probit Modelle); Markov-Ketten; Trendlinien; Spline-Kurven und anderes.
- ② Das *Modul Optimierung* wird zur Optimierung von mehrfachen Entscheidungsvariablen, die Einschränkungen unterstellt sind, verwendet, um ein Ziel zu maximieren oder minimieren. Das Modul kann folgendermaßen ausgeführt werden: entweder als eine statische Optimierung, oder als eine dynamische und stochastische Optimierung unter Ungewissheit zusammen mit einer Monte-Carlo-Simulation oder als eine stochastische Optimierung mit Super-Speed-Simulationen. Die Software kann lineare und nichtlineare Optimierungen mit binären, ganzzahligen und kontinuierlichen Variablen ausführen, sowie auch Markowitz effiziente Grenzen generieren.
- ② Das *Modul Analytische Tools* erlaubt Ihnen Folgendes auszuführen: Segmentierungsclustering, Hypothesentesten, statistische Tests von Rohdaten, Datendiagnosen von technischen Vorausberechnungshypothesen (z.B., Heteroskedastizität, Multikollinearität

und Ähnliches), Sensibilitäts- und Szenarioanalysen, Überlagerungsdiagrammanalysen, Spinnennetzdiagramme, Tornadodiagramme und viele andere leistungsstarke Tools.

- ② Der Super-Verband-Löser von Real Options ist eine weitere eigenständige Software, die den Risiko Simulator ergänzt und die verwendet wird, um einfache bis komplexe Probleme von Realoptionen zu lösen.

Die folgenden Abschnitte begleiten Sie durch die Grundlagen des Simulationsmoduls in Risiko Simulator, während spätere Kapitel Anwendungen von anderen Modulen in ausführlicherem Detail behandeln werden. Um mitzuverfolgen, vergewissern Sie sich, dass Risiko Simulator auf Ihrem Rechner installiert ist, bevor Sie fortfahren. *Eigentlich wird sehr empfohlen, dass Sie sich erst die Einleitungsvideos im Web (www.realoptionsvaluation.com/risksimulator.html) anschauen oder die schrittweisen Übungen am Ende dieses Kapitels versuchen, bevor Sie zur Exminierungs des in diesem Kapitel enthaltenen Texts zurückkehren.* Der Grund dafür liegt darin, dass sowohl die Videos als auch die Übungen Ihnen zum sofortigen Anfang helfen können, während der Text in diesem Kapitel mehr auf die Theorie und die detaillierten Erklärungen der Simulationseigenschaften gerichtet ist.

Eine Monte-Carlo-Simulation ausführen

Um eine Simulation in Ihrem existierenden Excelmodell auszuführen, müssen normalerweise die folgenden Schritte durchgeführt werden:

1. **Ein neues Simulationsprofil starten oder ein existierendes Profil öffnen**
2. **Inputhypothesen in den relevanten Zellen definieren**
3. **Outputvorausberechnungen in den relevanten Zellen definieren**
4. **Die Simulation ausführen**
5. **Die Ergebnisse interpretieren**

Wenn gewünscht, und zur Übung, öffnen Sie die Beispielsdatei genannt ***Grund-Simulationsmodell*** und befolgen Sie die nachstehenden Beispiele zur Kreierung einer Simulation. Die Beispielsdatei kann entweder vom Startmenü ***Start | Real Options Valuation | Risiko Simulator | Beispiele*** oder direkt durch ***Risk Simulator | Beispielsmodelle*** aufgerufen werden.

1. Ein neues Simulationsprofil starten

Um eine neue Simulation zu starten, müssen Sie erst ein neues Simulationsprofil erstellen. Ein Simulationsprofil enthält einen kompletten Satz von Anweisungen zur Ausführung Ihrer Simulation, das heißt, alle Hypothesen, Vausberechnungen, Laufpräferenzen und so weiter. Die Erstellung von Profilen erleichtert die Kreierung von mehrfachen Simulationsszenarien. Das heißt, unter Verwendung des exakt gleichen Modells kann man mehrere Profile kreieren, jedes

mit seinen eigenen spezifischen Simulationseigenschaften und -voraussetzungen. Dieselbe Person kann verschiedene Testszenarien unter Verwendung von Verteilungshypothesen und -inputs erstellen oder mehrere Personen können Ihre eigene Hypothesen und Inputs mit dem gleichen Modell testen.

- **Excel starten** und ein neues Modell kreieren oder ein existierendes Modell öffnen (sie können ein Grund-Simulation Beispielsmodell verwenden, um mitzuverfolgen)
- Klicken Sie auf **Risiko Simulator | Neues Simulationsprofil**
- Bestimmen Sie einen Titel für Ihre Simulation und spezifizieren Sie alle anderen relevanten Informationen (Bild 2.1)

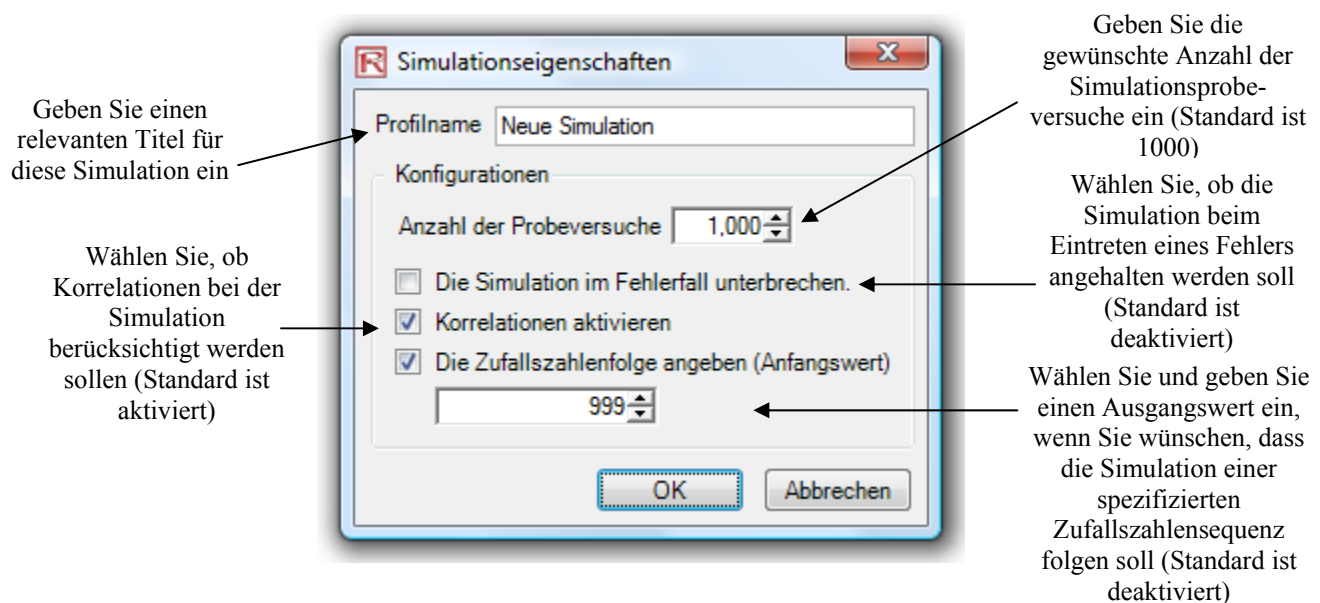


Bild 2.1—Neues Simulationsprofil

- ② **Titel:** Die Bestimmung eines Simulationstitels erlaubt Ihnen multiple Simulationsprofile in einem einzelnen Excelmodell zu kreieren. Dies bedeutet, dass Sie jetzt verschiedene Simulationsszenarienprofile innerhalb desselben Modells speichern können, ohne existierende Hypothesen löschen oder ändern zu müssen, wenn neue Simulationsszenarien erforderlich werden. Sie können den Profilnamen jederzeit ändern (**Risiko Simulator | Profil editieren**).
- ② **Anzahl der Probeversuche:** Hier geben Sie die erforderliche Anzahl der Simulationsprobeversuche ein. In anderen Worten, die Ausführung von 1000 Probeversuchen bedeutet, dass 1000 verschiedene Iterationen von Ausgängen, basierend auf den Inputhypothesen, generiert werden. Sie können dies beliebig ändern, aber die Eingabe muss aus positiven Ganzzahlen bestehen. Die Standardanzahl von Läufen ist 1000 Probeversuche. Sie können die Präzisions- und Fehlerkontrolle verwenden, um automatisch Hilfe bei der Bestimmung der auszuführenden Simulationsprobeversuche zu bekommen (siehe den Abschnitt über Präzisions- und Fehlerkontrolle für Details).

- ② **Die Simulation bei Fehler anhalten:** Wenn aktiviert, hält die Simulation bei jedem Auftreten eines Fehlers im Excelmodell an. Das heißt, wenn Ihr Modell auf einen Berechnungsfehler stößt (z.B., einige in einem Simulationsprobeversuch generierte Inputwerte könnten einen „teilen durch Null“ Fehler in einer Ihrer Tabellenblattzelle ergeben), wird die Simulation angehalten. Dies ist bei der Prüfung Ihres Modells wichtig, um sicherzugehen, dass sich keine Berechnungsfehler in Ihrem Excelmodell befinden. Wenn Sie allerdings über das Funktionieren Ihres Modells sicher sind, ist die Aktivierung dieser Präferenz nicht nötig.
- ② **Korrelationen aktivieren:** Wenn aktiviert, werden Korrelationen zwischen gepaarten Inputhypothesen berechnet. Andernfalls werden alle Korrelationen auf Null eingestellt und eine Simulation wird ausgeführt, unter der Annahme, dass es keine Kreuzkorrelationen zwischen den Inputhypothesen gibt. Als Beispiel, die Anwendung von Korrelationen liefert genauere Ergebnisse, wenn Korrelationen tatsächlich existieren, und tendiert eine geringere Vorausberechnungskonfidenz zu ergeben, wenn negative Korrelationen vorhanden sind. Nachdem Sie die Korrelationen hier aktiviert haben, können Sie später die relevanten Korrelationskoeffizienten für jede generierte Hypothese einstellen (siehe den Abschnitt über Korrelationen für mehr Details).
- ② **Zufallszahlensequenz bestimmen:** Eine Simulation wird definitionsgemäß leicht unterschiedliche Ergebnisse bei jeder Ausführung einer Simulation liefern. Dies geschieht Aufgrund der Zufallszahlen-generierungsroutine in einer Monte-Carlo-Simulation und ist ein theoretischer Fakt bei allen Zufallszahlengeneratoren. Wenn Sie allerdings einen Vortrag halten, benötigen Sie gelegentlich die gleichen Ergebnisse (insbesondere wenn der vorgeführte Bericht einen Ergebnissatz anzeigt und Sie möchten die Generierung der gleichen Ergebnisse während eines Live-Vortrags vorführen; oder wenn Sie Modelle mit anderen teilen und möchten die gleichen Ergebnisse jedes Mal erhalten). Dann, aktivieren Sie diese Einstellung und geben Sie eine anfängliche Ausgangszahl ein. Die Ausgangszahl kann eine beliebige positive Ganzzahl sein. Wenn Sie die gleiche anfängliche Ausgangszahl, die gleiche Anzahl von Probeversuchen und die gleiche Inputhypothesen verwenden, wird die Simulation immer die gleiche Sequenz von Zufallszahlen ergeben, was den gleichen Endergebnissatz gewährleistet.

Bitte bemerken Sie, dass wenn Sie ein neues Simulationsprofil kreiert haben, können Sie später zurückkommen und diese Auswahlen modifizieren. Um dies auszuführen, stellen Sie erst sicher, dass das aktuelle aktive Profil tatsächlich das gewünschte zu modifizierende Modell ist, sonst klicken Sie auf *Risiko Simulator | Simulationsprofil ändern*, wählen Sie das zu modifizierende Modell aus und klicken Sie auf **OK** (Bild 2.2 zeigt ein Beispiel mit mehrfachen Profilen an und wie man das ausgewählte Profil aktiviert). Dann klicken Sie auf *Risiko Simulator | Simulationsprofil editieren* und führen Sie die gewünschten Änderungen aus. Sie können auch ein existierendes Profil duplizieren oder umbenennen. Wenn Sie mehrfache Profile im gleichen Excelmodell erstellen, stellen Sie sicher, dass Sie jedem Profil einen einmaligen Namen geben, sodass Sie die Profile später auseinander halten können. Übrigens, da diese Profile innerhalb

versteckter Sektoren der Excel *.xls Datei gespeichert sind, müssen Sie keine zusätzlichen Dateien speichern. Die Profile und ihre Inhalte (Hypothesen, Vorausberechnungen und so weiter) werden automatisch gespeichert, wenn Sie die Exceldatei speichern. Zum Schluss, das letzte Profil, dass beim beenden und speichern der Exceldatei aktiv ist, wird auch das Profil sein, dass beim nächsten Zugriff auf der Exceldatei geöffnet sein wird.

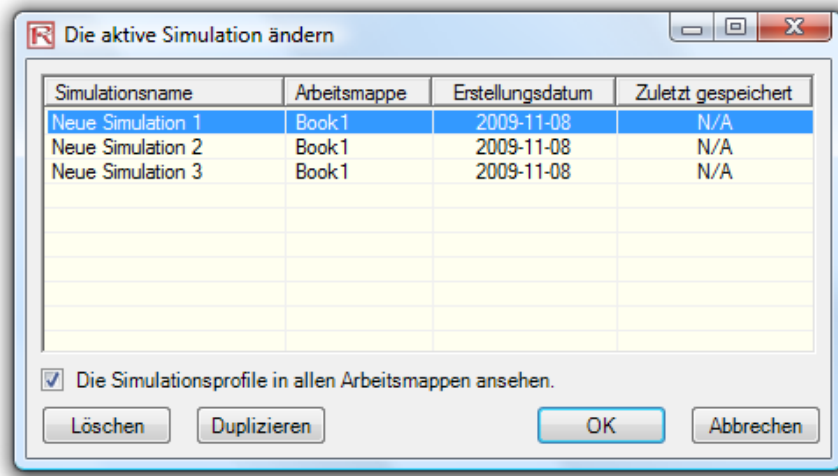


Bild 2.2—Aktive Simulation ändern

2. Inputhypothesen definieren

Der nächste Schritt ist die Inputhypothesen in Ihrem Modell festzulegen. Bitte bemerken Sie, dass man Hypothesen nur den Zellen zuordnen kann, die keine Gleichungen oder Funktionen (das heißt, eingegebene numerische Werte, welche Inputs in einen Modell sind) haben, während man Outputvorausberechnungen nur Zellen mit Gleichungen oder Funktionen (das heißt, Outputs eines Modells) zuordnen kann. Bitte gedenken Sie, dass man Hypothesen und Vorausberechnungen nicht einstellen kann, außer wenn ein Simulationsprofil schon existiert. Legen Sie neue Inputhypothesen in Ihrem Modell wie folgt fest:

- Stellen Sie sicher, dass ein Simulationsprofil existiert oder öffnen Sie ein existierendes Profil, oder starten Sie ein neues Profil (**Risiko Simulator | Neues Simulationsprofil**)
- Wählen Sie die Zelle, die einer Hypothese zugeordnet werden soll (z.B., Zelle **G8** im Beispiel des Grund-Simulationsmodells)
- Klicken Sie auf **Risiko Simulator | Inputhypothese einstellen** oder klicken Sie auf die Ikone Inputhypothese einstellen in der Ikonensymbolleiste von Risiko Simulator
- Wählen Sie die gewünschte relevante Verteilung, geben Sie die relevanten Verteilungsparameter (z.B., **Dreiecks**verteilung mit **1**, **2**, **2.5** als die minimalen, wahrscheinlichsten und maximalen Werte) ein und klicken Sie auf **OK**, um die Inputhypothese in Ihr Modell einzugeben (Bild 2.3)

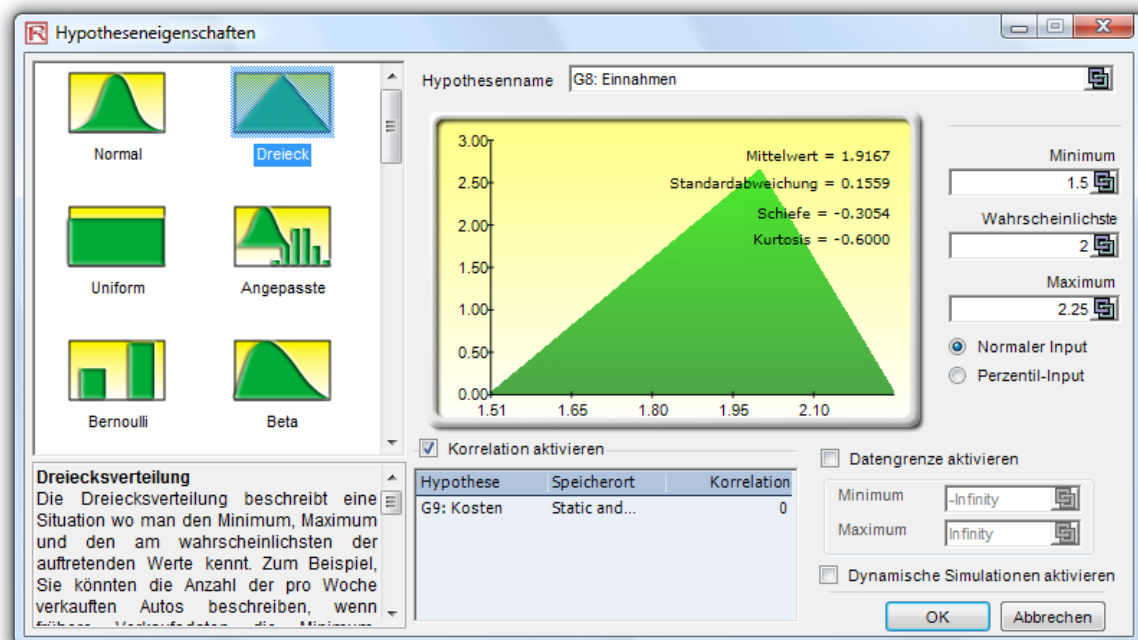


Bild 2.3—Eine Inputhypothese einstellen

Bitte bemerken Sie, dass Sie die Hypothesen auch folgendermaßen einstellen können. Wählen Sie die gewünschte Zelle aus, die den Hypothesen zuzuordnen ist. Rufen Sie das Kurzbefehlsmenü von Risiko Simulator unter Verwendung des **rechten Mausklicks** auf, um eine Inputhypothese einzustellen. Erfahrene Benutzer können außerdem Inputhypothesen unter Verwendung der **RS-Funktionen** von Risiko Simulator einstellen: Wählen Sie die gewünschte Zelle, klicken Sie auf *Einfügen, Funktion* in Excel, wählen Sie *Alle Kategorien* und scrollen Sie herunter zur Liste der RS-Funktionen (wir empfehlen die Verwendung der RS-Funktionen nicht, wenn sie kein erfahrener Benutzer sind). Für die folgenden Beispiele empfehlen wir, dass Sie die Grundanweisungen bei dem Zugriff auf Menüs und Ikonen befolgen.

Bitte bemerken Sie, dass es einige erwähnenswerte Schlüsselbereiche in den Hypotheseeseigenschaften gibt. Bild 2.4 zeigt die verschiedenen Bereiche an:

- ② **Hypothesenname:** Dies ist ein optionaler Bereich, wo Sie einmaligen Namen für die Hypothesen eingeben können, sodass Sie verfolgen können, was diese Hypothesen repräsentieren. Es ist gute Modellierungspraxis kurze aber präzise Hypothesennamen zu verwenden.
- ② **Verteilungsgalerie:** Dieser Bereich auf der linken Seite zeigt alle in der Software zur Verfügung stehenden Verteilungen an. Um die Ansichten zu ändern, klicken Sie rechts überall in der Galerie und wählen Sie große Ikonen, kleinen Ikonen oder Listen. Es stehen über zwei Dutzend Verteilungen zur Verfügung.
- ② **Die Inputparameter:** Die von der ausgewählten Verteilung abhängig erforderlichen relevanten Parameter werden hier angezeigt. Sie können die Parameter entweder direkt eingeben oder sie mit spezifischen Zellen in Ihrem Arbeitsblatt verknüpfen. Die

Hartkodierung oder manuelle Eingabe der Parameter ist nützlich, wenn man annimmt, dass die Hypothesenparameter sich nicht ändern. Die Verknüpfung mit Arbeitsblattzellen ist nützlich, wenn die Inputparameter sichtbar sein müssen oder sich ändern dürfen (klicken Sie auf die Verknüpfungssikone , um einen Inputparameter mit einer Arbeitsblattzelle zu verknüpfen.

- ② **Datengrenzen aktivieren:** Diese werden normalerweise nicht vom durchschnittlichen Analysten verwendet, sind aber vorhanden, um die Verteilungshypothesen zu stützen. Zum Beispiel, wenn eine Normalverteilung ausgewählt ist, befinden sich die theoretischen Grenzen zwischen der negativen Unendlichkeit und der positiven Unendlichkeit. In der Praxis, allerdings, existiert die simulierte Variable nur innerhalb eines kleineren Bereichs und man kann diesen Bereich eingeben, um die Verteilung passend zu stützen.
- ② **Korrelationen:** Hier kann man paarweise Korrelationen den Inputhypothesen zuordnen. Wenn Hypothesen erforderlich sind, denken Sie daran, die Präferenz *Korrelationen aktivieren* zu aktivieren, indem Sie auf [Risiko Simulator | Simulationsprofil editieren](#) klicken. Siehe die später in diesem Abschnitt vorgeführte Behandlung über Korrelationen für mehr Details bezüglich der Zuordnung von Korrelationen und deren Effekte auf ein Modell. Bitte bemerken Sie, dass Sie eine Verteilung entweder stützen oder sie mit einer anderen Hypothese korrelieren können, aber nicht beides.
- ② **Kurze Beschreibungen:** Diese sind für alle Verteilungen in der Galerie vorhanden. Die kurzen Beschreibungen erläutern sowohl wann die Anwendung einer bestimmten Verteilung angemessen ist als auch die Voraussetzungen für die Inputparameter. Siehe den Abschnitt in *Wahrscheinlichkeitsverteilungen für Monte-Carlo-Simulationen* begreifen für Details zu jeder in der Software verfügbaren Verteilung.
- ② **Normaler Input und Perzentil-Input:** Diese Option erlaubt dem Benutzer einen Sorgfaltsschnelltest der Inputhypothesen durchzuführen. Zum Beispiel, bei der Einstellung einer Normalverteilung mit bestimmten Inputs für den Mittelwert und die Standardabweichung, können Sie auf Perzentil-Input klicken, um die entsprechenden 10. und 90. Perzentile zu visualisieren.
- ② **Dynamische Simulation aktivieren:** Diese Option ist standardmäßig deaktiviert. Wenn Sie aber eine multidimensionale Simulation (d.h., wenn Sie die Inputparameter der Hypothese mit einer anderen Zelle, die selber eine Hypothese ist, verknüpfen, simulieren Sie die Inputs, oder simulieren Sie die Simulation) ausführen möchten, dann bedenken Sie diese Option zu aktivieren. Die dynamische Simulation wird nicht funktionieren, wenn die Inputs mit anderen wechselnden Inputhypothesen verknüpft sind.

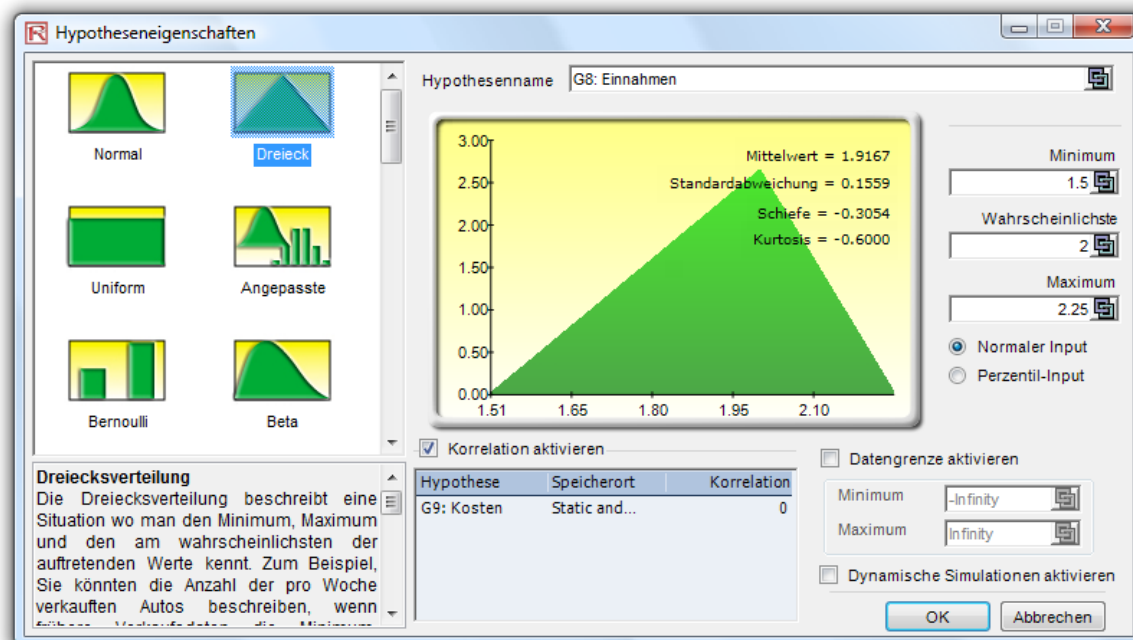


Bild 2.4—Hypotheseneigenschaften

Notiz: Wenn Sie dieses Beispiel verfolgen, fahren Sie fort mit der Einstellung einer anderen Hypothese auf Zelle **G9**. Verwenden Sie dieses Mal die **Uniform**verteilung mit einem Minimalwert von **0.9** und einem Maximalwert von **1.1**. Dann setzen Sie im nächsten Schritt mit der Definierung der Outputvorausberechnungen fort.

3. Outputvorausberechnungen definieren

Der nächste Schritt ist die Outputvorausberechnungen im Modell zu definieren. Vorausberechnungen (Prognosen) können nur auf Outputzellen mit Gleichungen oder Funktionen definiert werden. Folgendes beschreibt den Prozess der Einstellung einer Vorausberechnung:

- Wählen Sie die Zelle auf der Sie eine Hypothese einstellen möchten (z.B., Zelle **G10** in den Beispiel Simulationsgrundmodell)
- Klicken Sie auf **Risiko Simulator** und wählen Sie **Outputvorausberechnung einstellen** oder klicken Sie auf die Ikone Outputvorausberechnung einstellen in der Ikonen-Symbolleiste von Risiko Simulator (Bild 1.3)
- Geben Sie die relevanten Informationen ein und klicken Sie auf **OK**

Bitte bemerken Sie, dass Sie die Outputvorausberechnungen auch folgendermaßen einstellen können: Wählen Sie die Zelle auf der Sie eine Hypothese einstellen möchten und öffnen Sie das Kurzbefehlmenü von Risiko Simulator mit der rechten Maustaste, um eine Outputvorausberechnung einzustellen.

Das Bild 2.5 erläutert die Eigenschaften der Einstellung der Vorausberechnung.

- ② **Vorausberechnungsname:** Spezifizieren Sie den Namen der Vorausberechnungszelle. Dies ist wichtig weil, wenn Sie ein großes Modell mit mehrfachen

Vorausrechnungszellen haben, erlaubt Ihnen die individuelle Benennung der Vorausrechnungszellen einen schnellen Zugriff auf die richtigen Ergebnisse. Unterschätzen Sie nicht die Wichtigkeit dieses einfachen Schritts. Es ist gute Modellierungspraxis, kurze aber präzise Hypothesennamen zu verwenden.

- ② **Vorausrechnungspräzision:** Anstatt sich auf eine Daumenschätzung der Anzahl der in Ihrer Simulation auszuführenden Probeversuche verlassen zu müssen, können Sie Präzisions- und Fehlerkontrollen einstellen. Wenn eine Fehler-Präzision Kombination in der Simulation erreicht wurde, hält die Simulation an und informiert Sie über die Präzision. Dies macht die Anzahl der Simulationsprobeversuche zu einem automatisierten Prozess und bedarf so keiner Schätzungen der erforderlichen Anzahl der zu simulierenden Simulationsprobeversuche. Lesen Sie den Abschnitt über Fehler- und Präzisionskontrolle für ausführlichere Details.
- ② **Vorausrechnungsfenster anzeigen:** Erlaubt dem Benutzer ein bestimmtes Vorausrechnungsfenster anzuzeigen oder nicht. Die Standardeinstellung ist, das Vorausrechnungsdiagramm immer anzuzeigen.

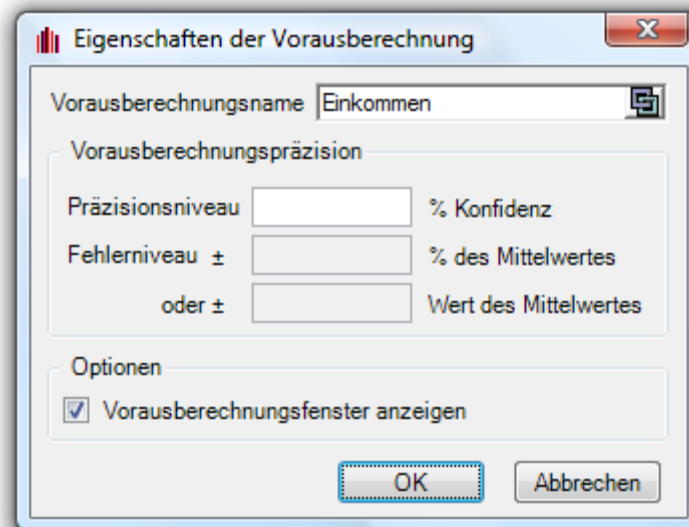


Bild 2.5—Outputvorausberechnung einstellen

4. Die Simulation ausführen

Wenn alles richtig ist, klicken Sie einfach auf **Risiko Simulator | Simulation ausführen** oder klicken Sie auf die Ikone **Ausführen** in der Symbolleiste von Risiko Simulator und die Simulation fährt fort. Sie können eine Simulation nach ihrem Lauf zurücksetzen, um sie wieder auszuführen (**Risiko Simulator | Simulation zurücksetzen** oder verwenden Sie die Ikone **Simulation zurücksetzen** in der Symbolleiste) oder sie während ihres Laufs anzuhalten. Außerdem, die Funktion **Schritt** (**Risiko Simulator | Schritt-Simulation** oder die Ikone **Schritt-Simulation** in die Symbolleiste) erlaubt Ihnen jeweils einen einzelnen Probeversuch zu

simulieren, nützlich wenn Sie andere über eine Simulation unterrichten (in anderen Worten, Sie können zeigen wie alle Werte in der Hypothesenzellen bei jeder Simulation ersetzt werden und das gesamte Modell jedes Mal neu berechnet wird). Sie haben auch Zugriff auf das Menü Simulation ausführen, wenn Sie irgendwo im Modell mit der rechten Maustaste klicken und Simulation ausführen wählen.

Risiko Simulator erlaubt Ihnen auch die Simulation mit äußerst hoher Geschwindigkeit auszuführen, genannt Super-Speed. Um dies auszuführen, klicken sie auf **Risiko Simulator | Super-Speed-Simulation ausführen** oder verwenden Sie die Ikone Super-Speed. Bemerken Sie wie viel schneller die Super-Speed-Simulation ausgeführt wird. Übungshalber führen Sie folgenden Vorgang aus: **Simulation zurücksetzen**, dann **Simulationsprofil editieren**, dann die **Anzahl der Probeversuche** auf **100000** ändern und **Super-Speed ausführen**. Die Ausführung sollte nur einige Sekunden dauern. Seien Sie sich allerdings bewusst, dass die Super-Speed-Simulation nicht ausgeführt werden kann, wenn das Modell Fehler, VBA (Visual Basic für Anwendungen) oder Verknüpfungen mit externen Datenquellen oder -anwendungen enthält. In solchen Situationen werden Sie benachrichtigt und die Simulation wird stattdessen bei normaler Geschwindigkeit ausgeführt. Simulation bei normaler Geschwindigkeit können trotz Fehler, VBA oder externen Verknüpfungen immer ausgeführt werden.

5. Die Vorausberechnungsergebnisse interpretieren

Der letzte Schritt in einer Monte-Carlo-Simulation ist, die resultierenden Vorausberechnungsdiagramme zu interpretieren. Bilder 2.6 bis 13 zeigen das Vorausberechnungsdiagramm und die entsprechenden nach Ausführung einer Simulation generierten Statistiken an. Normalerweise sind die folgenden Elemente bei der Interpretierung der Ergebnisse einer Simulation wichtig:

- ① **Vorausberechnungsdiagramm:** Das im Bild 2.6 angezeigte Vorausberechnungsdiagramm ist ein Wahrscheinlichkeits-histogramm, das die Frequenzanzahl der Werte anzeigt, die in der Gesamtanzahl der simulierten Probeversuche auftreten. Die vertikalen Balken zeigen die Frequenz des Auftretens eines bestimmten x Wertes aus der Gesamtanzahl der Probeversuche, während die kumulative Frequenz (glatte Linie), die Gesamtwahrscheinlichkeiten aller in der Vorausberechnung auftretenden Werten an und unter x anzeigen.
- ② **Vorausberechnungsstatistiken:** Die im Bild 2.7 angezeigten Vorausberechnungsstatistiken fassen die Verteilung der Vorausberechnungswerte zusammen, bezogen auf die vier Momente einer Verteilung. Siehe den Abschnitt *Die Vorausberechnungsstatistiken begreifen* für mehr Details über die Bedeutung einiger dieser Statistiken. Sie können zwischen den Leisten Histogramm und Statistiken wechseln, wenn Sie die Leertaste drücken.

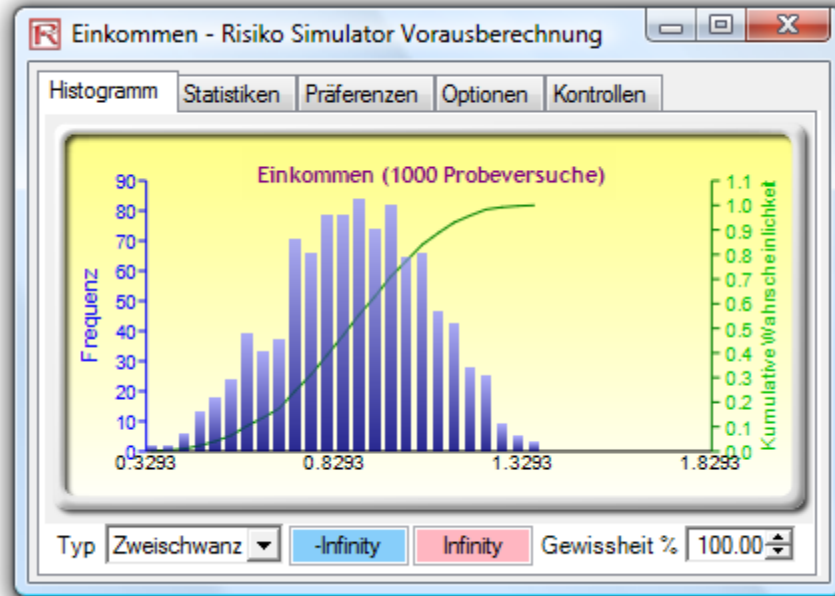


Bild 2.6—Vorausberechnungsdiagramm

Statistiken	Ergebnis
Anzahl der Probeversuche	1000
Mittelwert	0.8626
Medianwert	0.8674
Standardabweichung	0.1933
Varianz	0.0374
Variationskoeffizient	0.2241
Maximum	1.3570
Minimum	0.3019
Bereich	1.0551
Schiefe	-0.1157
Kurtosis	-0.4480
25% Perzentil	0.7269
75% Perzentil	1.0068
Prozent Fehlerpräzision bei 95% Konfidenz	1.3888%

Bild 2.7—Vorausberechnungsstatistiken

- ② **Präferenzen:** Die Leiste Präferenzen in dem Vorausberechnungsdiagramm erlaubt Ihnen das Erscheinungsbild der Diagramme zu ändern. Zum Beispiel, wenn *Immer im Vordergrund* ausgewählt ist, sind die Vorausberechnungsdiagramme immer sichtbar, unabhängig von der auf Ihrem Rechner laufenden Software. *Histogramm Auflösung* erlaubt Ihnen die Anzahl der Bins eines Histogramms zu ändern, von 5 Bins bis auf 100 Bins. Ferner, die Sektion *Daten aktualisieren* erlaubt Ihnen zu steuern, wie schnell die Simulation ausgeführt oder wie oft das Vorausberechnungsdiagramm aktualisiert wird. In

anderen Worten, wenn Sie wünschen, dass das Vorausberechnungsdiagramm fast bei jedem Probeversuch aktualisiert wird, wird das die Simulation verlangsamen, da mehr Speicher der Aktualisierung statt der Ausführung der Simulation zugeordnet ist. Dies ist lediglich eine Benutzerpräferenz und verändert keineswegs die Ergebnisse einer Simulation, nur die Geschwindigkeit ihrer Fertigstellung. Um die Geschwindigkeit der Simulation zusätzlich zu erhöhen, können Sie Excel während der Ausführung der Simulation minimieren. Dies reduziert den zur optischen Aktualisierung des Excelarbeitsblatts erforderlichen Speicher und macht ihn zur Ausführung der Simulation frei. *Alles schließen* und *Alles minimieren* steuern alle offenen Vorausberechnungsdiagramme.

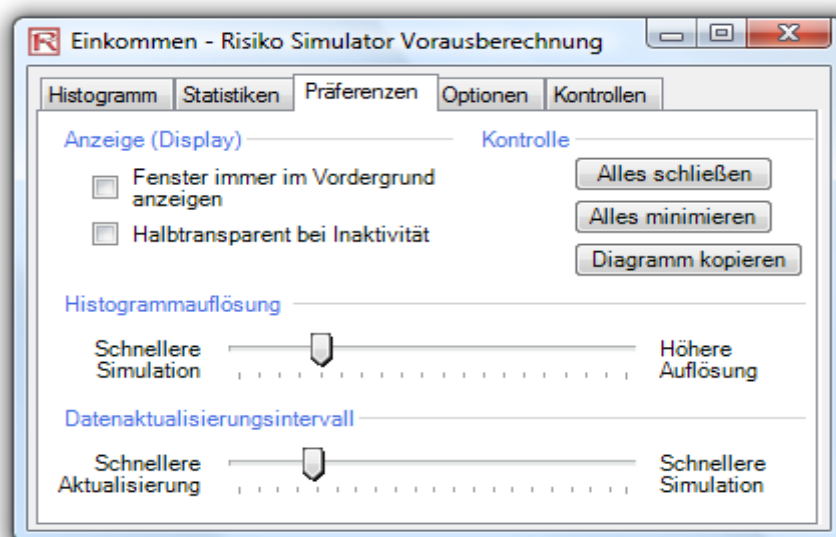


Bild 2.8—Vorausberechnungsdiagramm—Präferenzen

- ② **Optionen:** Diese Option des Vorausberechnungsdiagramm erlaubt Ihnen alle Vorausberechnungsdaten anzuzeigen oder nur einige Werte ein- oder auszufiltern, die innerhalb eines von Ihnen bestimmten Bereichs oder einer von Ihnen bestimmten Standardabweichung fallen. Sie können hier auch das Präzisionsniveau für diese spezifische Vorausberechnung einstellen, um die Fehlerniveaus in der Statistikansicht anzuzeigen. Siehe den Abschnitt über Fehler- und Präzisionskontrolle für mehr Details. *Die folgende Statistik anzeigen* ist eine Benutzerpräferenz: Sie bestimmen, ob die Linien des Mittelwerts, des Medianwerts, des 1. Quartils und des 4. Quartils (25. und 75. Perzentile) in dem Vorausberechnungsdiagramm angezeigt werden sollen.
- ② **Kontrolle:** Diese Leiste enthält alle Funktionalitäten, die Ihnen erlauben: Typ, Farbe, Größe, Zoom, Neigung, 3D und andere Elemente im Vorausberechnungsdiagramm zu ändern, Überlagerungsdiagramme (PDF, CDF) bereitzustellen, and Verteilungsanpassungen auf Ihre Vorausberechnungsdaten auszuführen (siehe die Abschnitte über Datenanpassung für mehr Details über diese Methodologie).

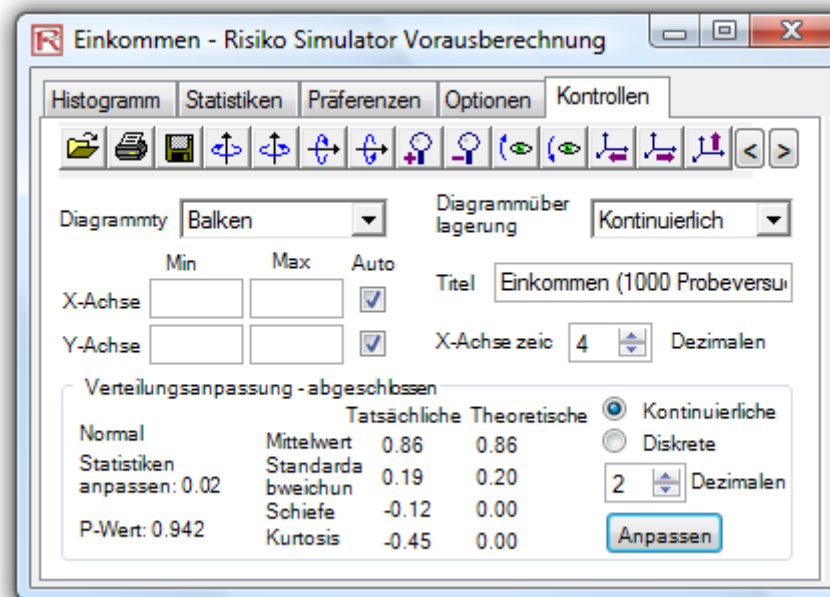
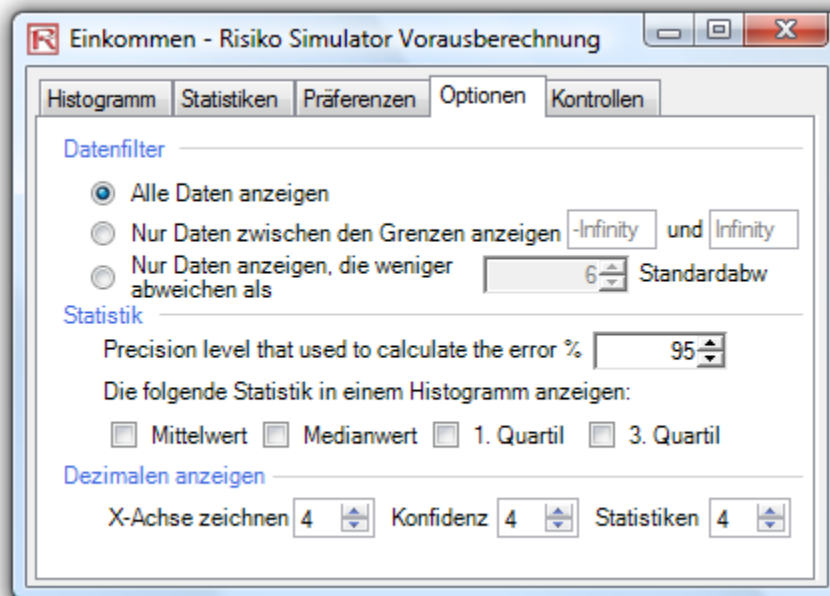


Bild 2.9—Vorausberechnungsdiagramm - Optionen und Kontrollen

Verwendung der Vorausberechnungsdiagramme und Konfidenzintervalle

In den Vorausberechnungsdiagrammen können Sie die Wahrscheinlichkeit des Auftretens, genannt Konfidenzintervalle, bestimmen. Das heißt, gegeben zwei Werte, was sind die Chancen, dass der Ausgang zwischen diese beiden Werte fallen wird? Bild 2.10 erläutert, dass es eine Wahrscheinlichkeit von 90% gibt, dass der Endausgang (in diesem Fall, das Niveau des Einkommens) zwischen \$0.2653 und \$1.3230 fallen wird. Das Zweisechwanz Konfidenzintervall kann folgendermaßen erhalten werden. Wählen Sie erst **Zweisechwanz** als den Typ, geben Sie

dann den gewünschten Gewissheitswert (z.B., **90**) ein und drücken Sie **TAB** auf der Tastatur. Die zwei berechneten Werte, die dem Gewissheitswert entsprechen, werden dann angezeigt. In diesem Beispiel, gibt es eine Wahrscheinlichkeit von 5%, dass das Einkommen unter \$0.2653 und eine weitere Wahrscheinlichkeit von 5%, dass das Einkommen über \$1.3230 sein wird. In anderen Worten, das Zweisechwanz Konfidenzintervall ist ein symmetrisches Intervall, dass um den Medianwert oder 50. Perzentilwert zentriert ist. Folglich haben die beiden Schwänze die gleiche Wahrscheinlichkeit.

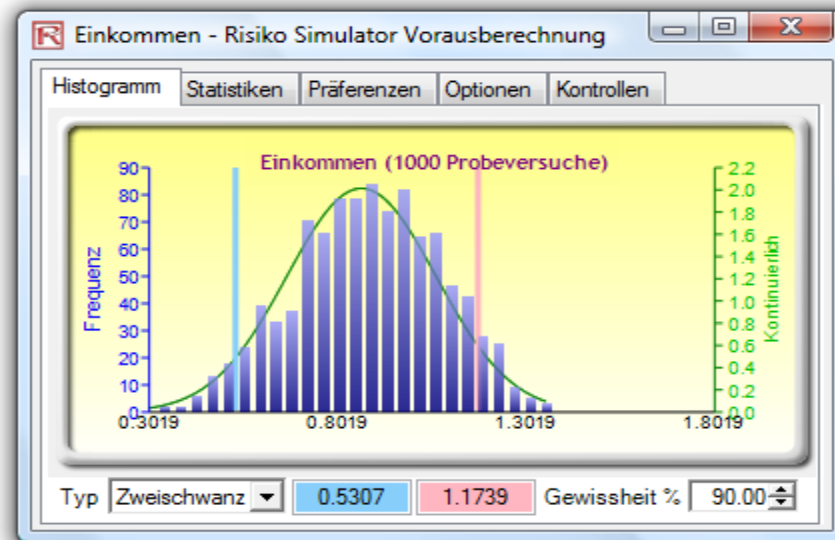


Bild 2.10—Vorausberechnungsdiagramm - Zweisechwanz Konfidenzintervall

Alternativ kann man eine Einschwanz Wahrscheinlichkeit berechnen. Bild 2.11 zeigt die Linksschwanz Auswahl bei einer Konfidenz von 95% (das heißt, wählen Sie **Linksschwanz** ≤ als den Typ, geben Sie **95** als das Gewissheitsniveau ein und drücken Sie auf **TAB** auf der Tastatur). Dies bedeutet, dass es eine Wahrscheinlichkeit von 95% gibt, dass das Einkommen unter \$1.3230, oder eine Wahrscheinlichkeit von 5% gibt, dass das Einkommen über \$1.3230 sein wird, was perfekt mit den im Bild 2.10 angezeigten Ergebnissen übereinstimmt.

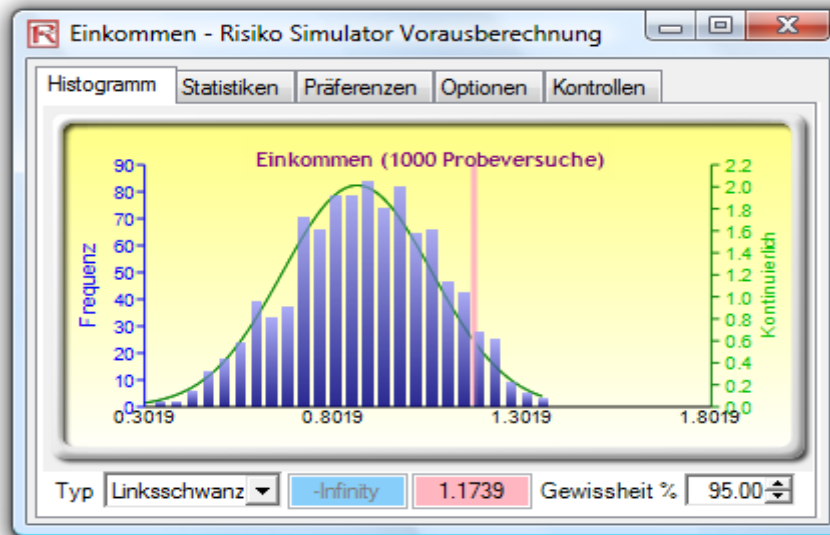


Bild 2.11—Vorausberechnungsdiagramm - Einschwanz Konfidenzintervall

Zusätzlich zur Bewertung des Konfidenzintervalls (d.h., gegeben ein Wahrscheinlichkeitsniveau, die entsprechenden Einkommenswerte zu finden), können Sie die Wahrscheinlichkeit eines gegebenen Einkommenswerts feststellen. Was ist, zum Beispiel, die Wahrscheinlichkeit, dass das Einkommen weniger als oder gleich \$1 sein wird? Um dies festzustellen, wählen Sie den Wahrscheinlichkeitstyp *Linksschwanz* $\leq$, geben Sie *1* im Inputfeld des Wertes ein und drücken Sie auf *TAB*. Es wird dann die entsprechende Gewissheit berechnet (in diesem Fall, gibt es Wahrscheinlichkeit von 67.70%, dass das Einkommen bei oder unter \$1 liegt).

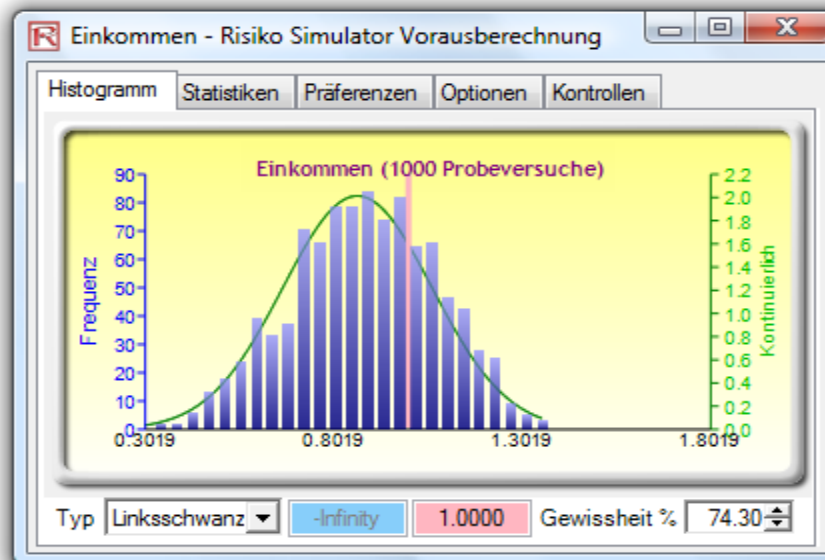


Bild 2.12—Vorausberechnungsdiagramm - Wahrscheinlichkeitsbewertung

Der Vollständigkeit halber können Sie den Wahrscheinlichkeitstyp *Rechtsschwanz* $>$ wählen, *1* im Inputfeld des Wertes eingeben und auf *TAB* drücken. Die resultierende Wahrscheinlichkeit

zeigt die Rechtsschwanz Wahrscheinlichkeit über den Wert von 1, das heißt, die Wahrscheinlichkeit eines Einkommens, das \$1 überschreitet (in diesem Fall sehen wir, dass es eine Wahrscheinlichkeit von 32.30% gibt, dass das Einkommen \$1 überschreitet). Die Summe von 67.70% und 32.30% ist natürlich 100%, die Gesamtwahrscheinlichkeit unter der Kurve.

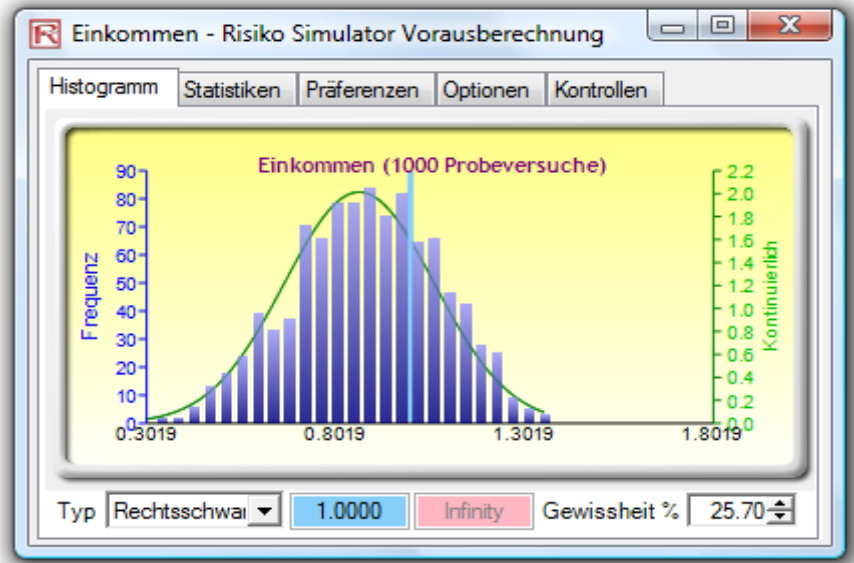


Bild 2.13—Vorausberechnungsdiagramm—Wahrscheinlichkeitsbewertung

TIPPS:

- Das Vorausberechnungsfenster ist größenveränderbar: Klicken Sie auf und ziehen Sie die untere rechte Ecke des Vorausberechnungsfensters. Es ist immer empfehlenswert, die aktuelle Simulation zurückzusetzen (*Risiko Simulator | Simulation zurücksetzen*), bevor Sie eine Simulation wieder ausführen.
- Denken Sie daran, dass Sie **TAB** auf der Tastatur drücken müssen, um das Diagramm und die Ergebnisse zu aktualisieren, wenn Sie die Gewissheits- oder Rechts- und Linksschwanzwerte eingeben.
- Sie können auch die **Leertaste** auf der Tastatur wiederholt drücken, um zwischen den verschiedenen Leisten (Histogramm, Statistiken, Präferenzen, Optionen und Kontrollen) zu wechseln.
- Außerdem, wenn Sie auf *Risiko Simulator | Optionen* klicken, bekommen Sie Zugang zu den verschiedenen Optionen von Risiko Simulator. Sie können Folgendes ausführen: Risiko Simulator automatisch bei jedem Start von Excel mitstarten oder nur wenn Sie es wünschen (gehen Sie zu *Start | Programme | Real Options Valuation | Risiko Simulator | Risiko Simulator*); die **Zellenfarben** der Hypothesen und Vorausberechnungen ändern; die **Zellenkommentare** ein- und ausschalten (Zellenkommentare erlauben Ihnen zu unterscheiden, welche Zellen Inputhypothesen und welche Outputvorausberechnungen sind, und ihre jeweiligen Inputparameter und Namen zu visualisieren). Verbringen Sie einige Zeit beim probieren der Outputs und Funktionalitäten, insbesondere die Leiste **Kontrollen**, der Vorausberechnungsdiagramme.

Die Grundlagen der Korrelationen

Der Korrelationskoeffizient ist eine Maßeinheit der Schärfe und Richtung des Verhältnisses zwischen zwei Variablen und kann alle Werte zwischen -1.0 und $+1.0$ annehmen. Das heißt, der Korrelationskoeffizient kann in seinem Signum (positives oder negatives Verhältnis zwischen zwei Variablen) und der Größe oder Schärfe des Verhältnisses (je höher der absolute Wert des Korrelationskoeffizienten, umso stärker das Verhältnis) zerlegt werden.

Der Korrelationskoeffizient kann auf verschiedene Weise berechnet werden. Die erste Methode ist, die Korrelation r von zwei Variablen x und y unter Verwendung folgender Formel manuell zu berechnen:

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

Die zweite Methode ist, die Funktion *CORREL* von Excel zu verwenden. Zum Beispiel, wenn die 10 Datenpunkte für x und y in den Zellen A1:B10 aufgelistet sind, dann ist die zu verwendende Excelfunktion *CORREL (A1:A10, B1:B10)*.

Die dritte Methode ist, das *Multi-Anpassungs-Tool* von *Risiko Simulator* auszuführen: Die resultierende Korrelationsmatrix wird berechnet und visualisiert.

Es ist wichtig zu bemerken, dass die Korrelation keine Kausalität impliziert. Zwei völlig unverwandte Zufallsvariablen können irgendeine Korrelation vorweisen, aber dies impliziert keine Kausalität zwischen den beiden (z.B., Sonnenfleckenaktivität und Ereignisse im Aktienmarkt sind zwar korreliert, aber es besteht keine Kausalität zwischen den beiden).

Es gibt zwei allgemeine Arten von Korrelationen: parametrische and nicht-parametrische Korrelationen. Der Pearsonsche Korrelationskoeffizient ist die geläufigste Korrelationsmaßeinheit und wird normalerweise einfach als der Korrelationskoeffizient bezeichnet. Allerdings ist die Pearsonsche Korrelation eine parametrische Maßeinheit, was bedeutet, dass sie Folgendes erfordert: beide korrelierten Variablen müssen eine unterliegende Normalverteilung haben und das Verhältnis zwischen den Variablen ist linear. Wenn diese Bedingungen verletzt werden, was häufig bei Monte-Carlo-Simulationen vorkommt, werden die nicht-parametrischen Gegenstücke wichtiger. Die Spearmansche Rangkorrelation and Kendall's Tau sind zwei Alternativen. Die Spearman-Korrelation wird am häufigsten verwendet und ist meist geeignet, wenn im Rahmen einer Monte-Carlo-Simulation angewendet — es gibt keine Abhängigkeit bei Normalverteilungen oder Linearität, was bedeutet, dass man Korrelationen

zwischen verschiedenen Variablen mit unterschiedlicher Verteilung anwenden kann. Um die Spearman-Korrelation zu berechnen, arrangieren Sie erst alle Werte der x und y Variablen und wenden Sie dann die Pearsonsche Korrelationsberechnung an.

Im Falle des Risiko Simulators, ist die verwendete Korrelation die mehr robuste nicht-parametrische Spearmansche Rangkorrelation. Allerdings, um den Simulationsprozess zu vereinfachen und mit der Korrelationsfunktion von Excel's überein zu stimmen, sind die erforderlichen Korrelationsinputs der Pearsonsche Korrelationskoeffizient. Risiko Simulator wird dann seine eigenen Algorithmen anwenden, um sie in die Spearmansche Rangkorrelation zu konvertieren, was den Prozess vereinfacht. Um jedoch die Benutzeroberfläche zu vereinfachen, erlauben wir den Benutzern die geläufigere Pearsonsche Produkt-Moment-Korrelation einzugeben (z.B., berechnet unter Verwendung der Funktion *CORREL* von Excel), während in den mathematischen Codes, konvertieren wir diese einfache Korrelationen in Spearmansche Rangkorrelationen für Verteilungssimulationen.

Anwendung von Korrelationen in Risiko Simulator

Man kann Korrelationen in Risiko Simulator auf mehrere weisen anwenden:

- ② Wenn man Hypothesen definiert (***Risiko Simulator | Inputhypothese einstellen***), die Korrelationen einfach in den Korrelationsmatrixraster in der Verteilungen-Galerie eingeben.
- ② Bei existierenden Daten, benutzen Sie das Multi-Anpassungs-Tool (***Risiko Simulator | Tools | Verteilungsanpassung | Merfache Variablen***), um eine Verteilungsanpassung durchzuführen und die Korrelationsmatrix zwischen paarweisen Variablen zu erhalten. Wenn ein Simulationsprofil vorhanden ist, werden die angepassten Hypothesen automatisch die relevanten Korrelationswerte enthalten.
- ② Bei existierenden Hypothesen, können Sie auf ***Risiko Simulator | Tools | Korrelationen editieren*** klicken, um die paarweisen Korrelationen aller Hypothesen direkt in eine Benutzeroberfläche einzugeben.

Bitte bemerken Sie, dass die Korrelationsmatrix positiv definit sein muss. Das heißt, dass die Korrelation mathematisch gültig sein muss. Zum Beispiel, angenommen Sie versuchen drei Variable zu korrelieren: Zensuren von Hochschulstudenten in einem bestimmten Jahr, die Anzahl der von ihnen konsumierten Biere pro Woche und die Anzahl ihrer dem Studieren gewidmeten Stunden pro Woche. Man würde annehmen, dass die folgenden Korrelationsverhältnisse existieren:

Zensuren und Bier: –	Je mehr sie trinken, umso niedriger die Zensuren (Nichtantritt bei Examen)
Zensuren und Studieren: +	Je mehr sie studieren, umso höher die Zensuren
Bier und Studieren: –	Je mehr sie trinken, umso weniger studieren sie (ständig betrunken und feiernd)

Allerdings, wenn Sie eine negative Korrelation zwischen Zensuren und Studieren eingeben und annehmen, dass die Korrelationskoeffizienten hohe Größen haben, wird die Korrelationsmatrix nichtpositiv definit sein. Das würde der Logik, den Korrelationsvoraussetzungen und der Matrixmathematik widersprechen. Kleinere Koeffizienten jedoch können trotz der schlechten Logik gelegentlich funktionieren. Wenn Sie eine nichtpositive oder schlechte Korrelationsmatrix eingeben, wird Risiko Simulator Sie automatisch informieren und anbieten, die Korrelationen in etwas semi-positiv definit zu korrigieren, beibehaltend dabei die Gesamtstruktur des Korrelationsverhältnisses (die gleiche Signa sowie auch die gleichen relativen Schärfen).

Die Effekte von Korrelationen in Monte-Carlo-Simulationen

Obwohl die bei der Korrelierung von Variablen in eine Simulation erforderlichen Berechnungen komplex sind, sind die resultierenden Ergebnisse ziemlich eindeutig. Bild 2.14 zeigt ein einfaches Korrelationsmodell (Korrelationseffekte Modell im Ordner Beispiele). Die Berechnung für Einnahmen ist einfach Preis multipliziert mit Menge. Dasselbe Modell wird für keine Korrelationen, positive Korrelationen (+0.9) und negative Korrelationen (−0.9) zwischen Preis und Menge wiederholt.

	Korrelationsmodell		
	Ohne Korrelation	Positive Korrelation	Negative Korrelation
Preis	\$2.00	\$2.00	\$2.00
Menge	1.00	1.00	1.00
Einnahmen	\$2.00	\$2.00	\$2.00

*Um dieses Modell zu replizieren, verwenden Sie die folgenden Hypothesen:
 Preise sind als Dreiecksverteilungen (1.8, 2.0, 2.2) eingestellt, während
 Menge als Uniformverteilungen (0.9, 1.1) eingestellt sind mit den Korrelationen
 auf 0,0, +0,8, −0,8 eingestellt bei 1000 Probeversuchen mit Ausgangswert 123456.*

Bild 2.14—Einfaches Korrelationsmodell

Die resultierenden Statistiken werden im Bild 2.15 angezeigt. Bitte bemerken sie, dass die Standardabweichung des Modells ohne Korrelationen 0.1450 ist, verglichen mit 0.1886 für die positive Korrelation und 0.0717 für die negative Korrelation. Das heißt, bei einfachen Modellen tendieren negative Korrelationen, die durchschnittliche Dispersion der Verteilung zu reduzieren and eine dichte und konzentriertere Vorausberechnungsverteilung zu kreieren, im Vergleich zu positiven Korrelationen mit größeren durchschnittlichen Dispersionen. Der Mittelwert jedoch bleibt relativ stabil. Dies impliziert, dass Korrelationen wenig mit der Änderung der Erwartungswerte von Projekten zu tun haben, aber das Risiko eines Projektes verringern oder vergrößern können.

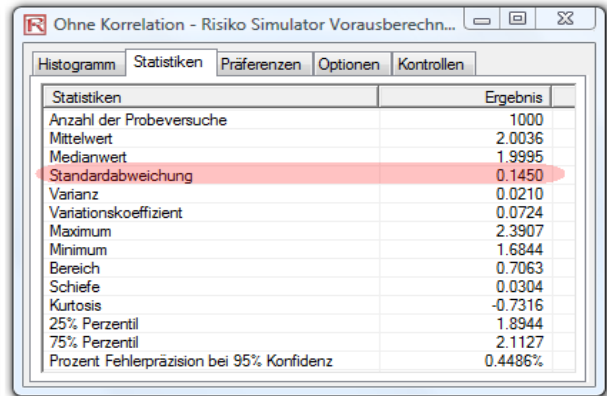
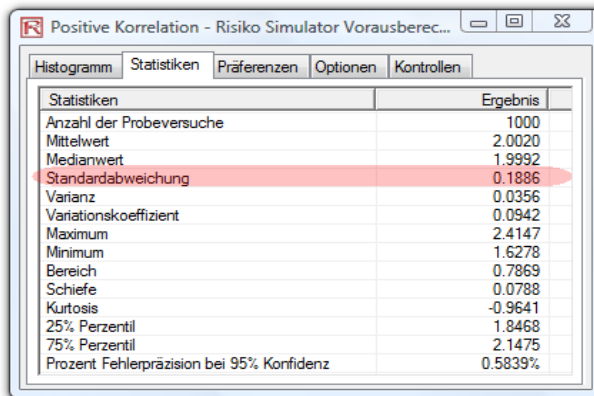
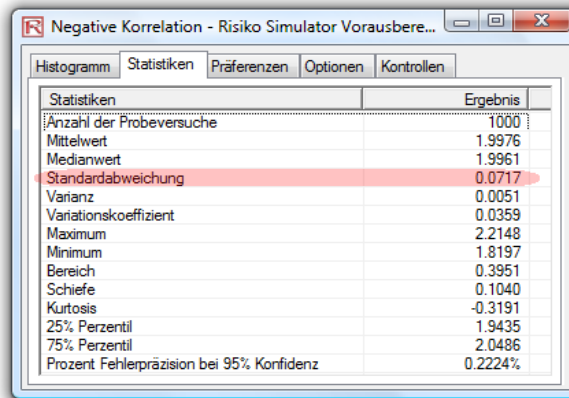


Bild 2.15—Ergebnisse von Korrelationen

Bild 2.16 stellt die Ergebnisse nach Ausführung einer Simulation, Extrahierung der Rohdaten der Hypothesen und Berechnung der Korrelationen zwischen den Variablen dar. Das Bild zeigt, dass die Inputhypothesen in der Simulation wiederhergestellt sind. Das heißt, Sie geben +0.9 und -0.9 Korrelationen ein und die resultierenden simulierten Werte haben die gleichen Korrelationen.

<i>Preis</i>	<i>Menge</i>		<i>Preis</i>	<i>Menge</i>	
<i>Positive</i>	<i>Positive</i>		<i>Negative</i>	<i>Negative</i>	
<i>Korrelation</i>	<i>Korrelation</i>		<i>Korrelation</i>	<i>Korrelation</i>	
1.95	0.91		1.89	1.06	
1.92	0.95		1.98	1.05	
2.02	1.04	Pearson-Korrelation:	1.89	1.09	Pearson-Korrelation:
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	
1.87	0.91		1.86	1.04	
1.84	0.91		1.96	1.02	
2.06	1.03		1.90	1.02	
1.98	1.01		1.92	1.10	

Bild 2.16—Wiederhergestellte Korrelationen

Präzisions- und Fehlerkontrolle

Ein sehr leistungsstarkes Tool in Monte-Carlo-Simulationen ist die Präzisionskontrolle. Zum Beispiel, wie viele auszuführenden Probeversuche gelten als ausreichend in einem komplexen Modell? Die Präzisionskontrolle vermeidet die Spekulation bei der Schätzung der relevanten Probeversuche, indem Sie der Simulation erlaubt anzuhalten, wenn das vorgegebene Präzisionsniveau erreicht wird.

Die Präzisionskontrollfunktionalität erlaubt es Ihnen zu bestimmen, wie präzise Ihre Vorausberechnung sein sollte. Allgemein betrachtet, je mehr Probeversuche berechnet werden, desto schmaler wird das Konfidenzintervall und umso genauer werden die Statistiken. Die Präzisionskontrollfunktion in Risiko Simulator verwendet die Eigenschaften der Konfidenzintervalle, um festzustellen, wann die spezifizierte Genauigkeit einer Statistik erreicht wurde. Für jede Vorausberechnung können Sie das spezifische Konfidenzintervall für das Präzisionsniveau bestimmen.

Achten Sie darauf, dass Sie diese sehr unterschiedlichen Begriffe nicht verwechseln: Fehler, Präzision und Konfidenz. Obwohl sie sich ähnlich anhören, sind die Konzepte erheblich verschieden. Ein einfaches Beispiel ist angebracht. Nehmen wir an, dass Sie ein Taco-Shell-Hersteller sind. Sie sind daran interessiert herauszufinden, wie viele gebrochene Taco-Shells es durchschnittlich in einer Packung von 100 Shells gibt. Eine Art dies herauszufinden ist, eine Stichprobe von abgepackten Packungen von 100 Taco-Shells zu sammeln, diese zu öffnen und die Anzahl der tatsächlich gebrochenen Taco-Shells zu zählen. Sie stellen 1 Million Packungen pro Tag her (das ist Ihre *Bevölkerung*), aber Sie öffnen stichprobenartig nur 10 Packungen (das ist Ihre *Stichprobengröße*, auch als die Anzahl Ihrer *Probenversuche* in einer Simulation) bekannt. Die Anzahl der gebrochenen Shells in jeder Packung ist wie folgt: 24, 22, 4, 15, 33, 32, 4, 1, 45 und 2. Die berechnete Durchschnittsanzahl der gebrochenen Shells ist 18.2. Basierend auf diesen 10 Stichproben oder Probeversuchen, liegt der Durchschnitt bei 18.2 Einheiten; basierend indes auf der Stichprobe, liegt das 80% Konfidenzintervall zwischen 2 und 33 Einheiten (das heißt, 80% der Male, liegt die Anzahl der gebrochenen Shells zwischen 2 und 33, *basierend auf dieser Stichprobengröße oder Anzahl der Probenversuchläufe*). Aber wie sicher sind Sie, dass 18.2 der richtige Durchschnittswert ist? Genügen 10 Probeversuche, um dies festzustellen? Das Konfidenzintervall zwischen 2 und 33 ist zu breit und zu variabel. Nehmen wir an, Sie benötigen einen genaueren Durchschnittswert, wobei der Fehler bei ± 2 Taco-Shells in 90% der Male liegt—dies bedeutet, dass wenn Sie *alle* 1 Million pro Tag hergestellten Packungen öffnen, 900000 dieser Packungen werden gebrochene Taco-Shells durchschnittlich bei einer Mittelwerteinheit ± 2 Tacos enthalten. Wie viele weitere Taco-Shell-Packungen müssten Sie dann probieren (bzw. Probeversuche ausführen), um dieses Präzisionsniveau zu erhalten? Hier sind die 2 Tacos das Fehlerniveau, während 90% das Präzisionsniveau ist. Wenn eine ausreichende Anzahl von Probeversuchen ausgeführt wird, dann wird das 90% Konfidenzintervall identisch mit dem 90% Präzisionsniveau sein, wobei ein genaueres Maß des Durchschnitts erhalten wird, sodass 90% der

Male, der Fehler und, deshalb, die Konfidenz ± 2 Tacos sein wird. Als Beispiel, sagen wir mal, dass der Durchschnitt 20 Einheiten ist, dann wird das 90% Konfidenzintervall zwischen 18 und 22 Einheiten liegen. Dieses Intervall ist präzise 90% der Male, wobei wenn man alle 1 Million Packungen öffnet, werden 900000 davon zwischen 18 und 22 gebrochene Tacos enthalten. Die Anzahl der erforderlichen Probeversuche, um diese Präzision zu treffen basiert auf der Stichprobenfehlergleichung $\bar{x} \pm Z \frac{s}{\sqrt{n}}$. Dabei ist $Z \frac{s}{\sqrt{n}}$ der Fehler von 2 Tacos, $\bar{x}$ ist der Stichprobendurchschnitt, Z ist der Standard-Normal Z-Wert, der von dem 90% Präzisionsniveau erhalten wird, s ist die Stichproben-Standardabweichung und n ist die Anzahl der erforderlichen Probeversuche, um dieses Fehlerniveau mit der spezifizierten Präzision zu treffen. Die Bilder 2.17 und 2.18 zeigen wie man die Präzisionskontrolle auf mehrfachen simulierten Vorausberechnungen in Risiko Simulator ausführen kann. Mit dieser Eigenschaft muss der Benutzer nicht mehr entscheiden, wie viele Probeversuche in einer Simulation auszuführen sind; alle Spekulationen werden eliminiert. Bild 2.17 zeigt ein Vorausberechnungsdiagramm mit einem bei 95% eingestelltem Präzisionsniveau. Dieser Wert kann verändert werden und wird in der Leiste Statistiken reflektiert, wie im Bild 2.18 angezeigt.

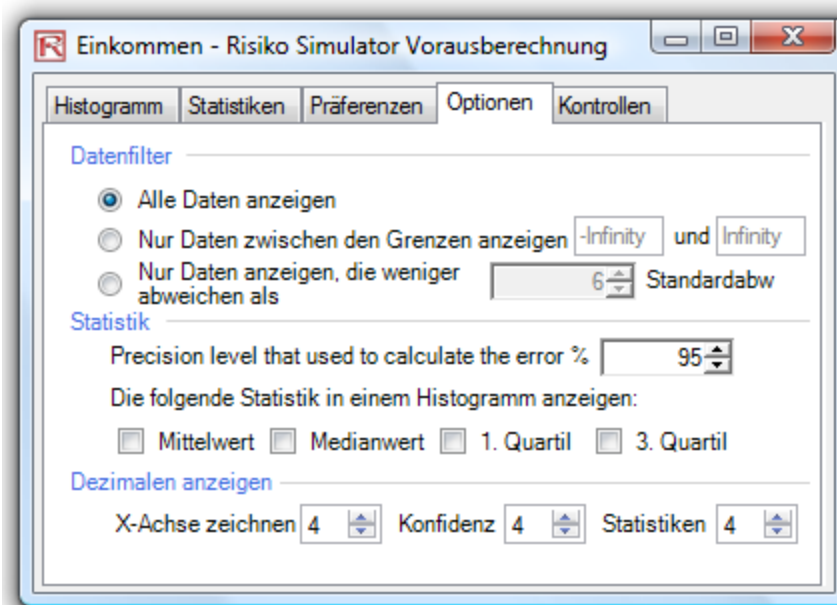
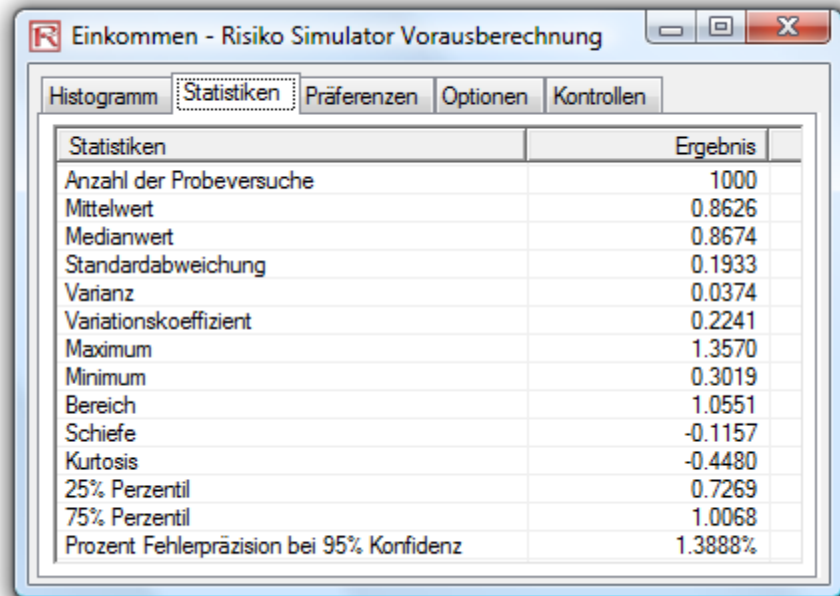


Bild 2.17—Einstellung des Präzisionsniveaus der Vorausberechnung



The screenshot shows a software window titled "Einkommen - Risiko Simulator Vorausberechnung". It has a tabbed interface with five tabs: "Histogramm", "Statistiken", "Präferenzen", "Optionen", and "Kontrollen". The "Statistiken" tab is currently selected. Below the tabs is a table with two columns: "Statistiken" and "Ergebnis". The table contains the following data:

Statistiken	Ergebnis
Anzahl der Probeversuche	1000
Mittelwert	0.8626
Medianwert	0.8674
Standardabweichung	0.1933
Varianz	0.0374
Variationskoeffizient	0.2241
Maximum	1.3570
Minimum	0.3019
Bereich	1.0551
Schiefte	-0.1157
Kurtosis	-0.4480
25% Perzentil	0.7269
75% Perzentil	1.0068
Prozent Fehlerpräzision bei 95% Konfidenz	1.3888%

Bild 2.18—Berechnung des Fehlers

Die Vorausberechnungsstatistiken begreifen

Die meisten Verteilungen können innerhalb von vier Momenten beschrieben werden. Das erste Moment beschreibt ihre Position oder Zentraltendenz (Erwartungserträge), das zweite Moment beschreibt ihre Breite oder Dispersion (Risiken), das dritte Moment beschreibt ihre Richtungsschiefe (die wahrscheinlichsten Ereignisse) und das vierte Moment beschreibt ihre Spitzigkeit oder die Dicke ihrer Schwänze (katastrophale Verluste oder Gewinne). Man sollte alle vier Momente in der Praxis berechnen und interpretieren, um eine umfassendere Sicht des unter Analyse stehenden Projekts zu bekommen. Risiko Simulator liefert die Ergebnisse aller vier Momente in der Leiste *Statistiken* in den Vorausberechnungsdiagramme.

Vermessen des Verteilungszentrums — das erste Moment

Das erste Moment einer Verteilung misst die erwartete Ertragsrate eines bestimmten Projekts. Es misst die Position der Projektszenarien und die möglichen durchschnittlichen Ereignisse. Die übliche Statistiken für das erste Moment schließen den Mittelwert (Durchschnitt), den Medianwert (Verteilungszentrum) und den Modalwert (den am häufigsten vorkommenden Wert) ein. Bild 2.19 stellt das erste Moment dar — wobei, in diesem Fall, wird der erste Moment dieser Verteilung durch seinen Mittelwert (μ) oder Durchschnittswert gemessen.

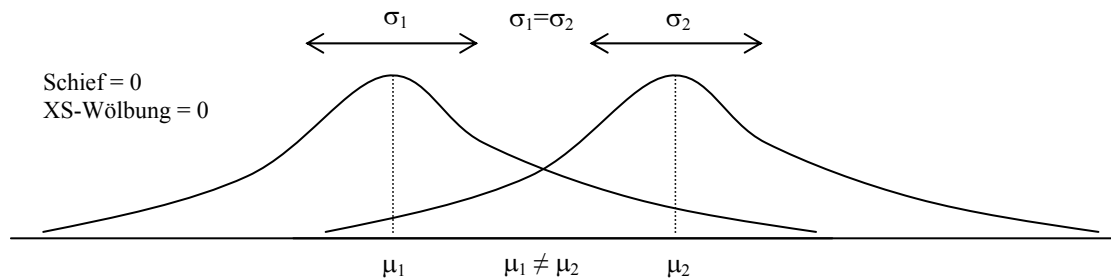


Bild 2.19—Erstes Moment

Vermessen der Verteilungsdispersion — das zweite Moment

Das zweite Moment misst die Dispersion einer Verteilung, welche eine Maßeinheit des Risikos ist. Die Dispersion oder Breite einer Verteilung misst die Variabilität einer Variablen, das heißt, das Potential, dass die Variable in unterschiedlichen Bereiche der Verteilung fallen kann— in anderen Worten, die potentiellen Ausgangsszenarien. Bild 2.20 zeigt zwei Verteilungen mit identischen ersten Momenten (identische Mittelwerte) aber sehr unterschiedlichen zweiten Momenten oder Risiken. Die Visualisierung wird im Bild 2.21 deutlicher. Als Beispiel, nehmen wir an, dass es zwei Aktien und die Bewegungen der ersten Aktie (dargestellt durch die dunkleren Linie) mit den kleineren Fluktuationen werden gegen die Bewegungen der zweiten Aktie (dargestellt durch die punktierte Linie) mit einer viel höheren Preisfluktuation verglichen.

Offensichtlich würde ein Investor die Aktie mit der heftigeren Fluktuation als risikoreicher betrachten, weil die Ausgänge der risikoreicheren Aktie relativ unbekannter als die der weniger risikoreicheren Aktie sind. Die Vertikalachse im Bild 2.21 misst die Aktienpreise, deshalb hat die risikoreichere Aktie einen breiteren Bereich von potentiellen Ausgängen. Dieser Bereich wird in die Breite der Verteilung (die Horizontalachse) im Bild 2.20 konvertiert, wobei die breitere Verteilung das risikoreichere Aktivum repräsentiert. Daher misst die Breite oder Dispersion einer Verteilung die Risiken einer Variablen.

Mit Bezug auf das Bild 2.20, bitte bemerken Sie, dass obwohl beide Verteilungen identische erste Momente oder Zentraltendenzen haben, sind die Verteilungen offensichtlich sehr unterschiedlich. Dieser Unterschied in der Verteilungsbreite ist messbar. Mathematisch und statistisch kann man die Breite oder das Risiko einer Variablen mittels mehrerer verschiedener Statistiken messen, einschließlich Bereich, Standardabweichung (σ), Varianz, Variationskoeffizient und Perzentile.

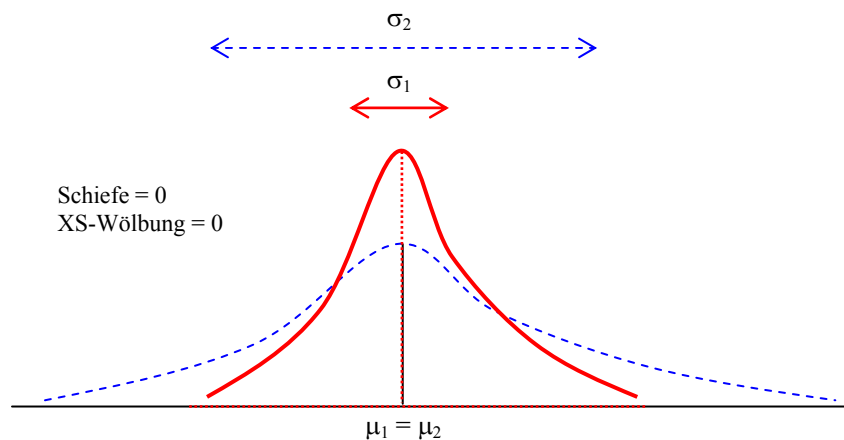


Bild 2.20—Zweites Moment

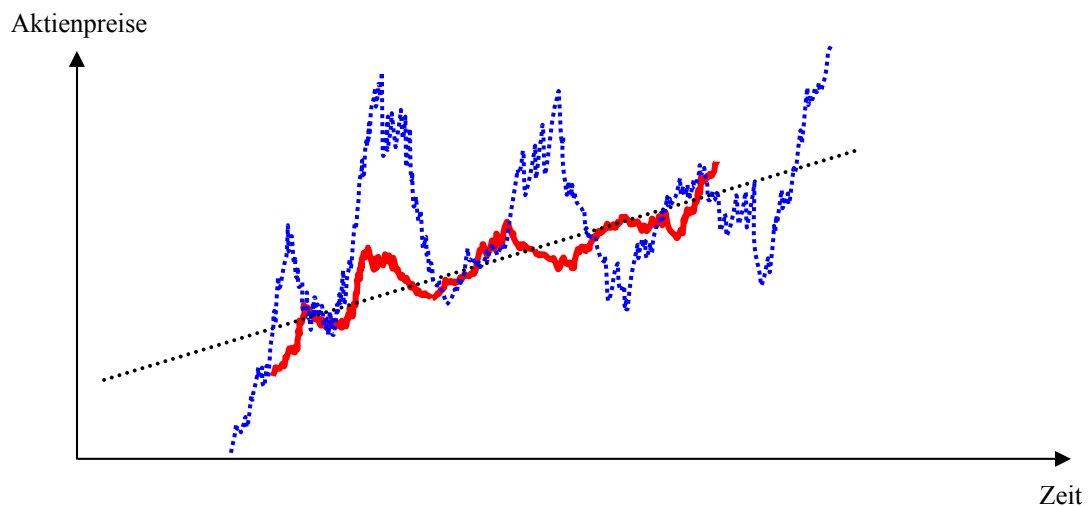


Bild 2.21 –Fluktuationen der Aktienpreise

Vermessen der Verteilungsschiefe — das dritte Moment

Das dritte Moment misst die Schiefe einer Verteilung, in anderen Worten, wie die Verteilung zu der einen oder der anderen Seite gezogen wird. Bild 2.22 stelle eine negative oder linke Schiefe (der Schwanz der Verteilung zeigt nach links) und Bild 2.23 stelle eine positive oder rechte Schiefe (der Schwanz der Verteilung zeigt nach rechts) dar. Der Mittelwert ist immer in der Richtung des Verteilungsschwanzes verzerrt, während der Medianwert konstant bleibt. Eine andere Weise dies anzusehen ist, dass der Mittelwert sich bewegt, aber die Standardabweichung, die Varianz oder die Breite können konstant bleiben. Wenn das dritte Moment nicht berücksichtigt wird und man sich nur auf die Erwartungserträge (z.B., Medianwert oder Mittelwert) und das Risiko (Standardabweichung) bezieht, könnte sich die Wahl eines Projekts mit positiver Schiefe als fehlerhaft erweisen! Zum Beispiel, wenn die Horizontalachse die Nettoeinnahmen eines Projekts repräsentiert, dann wäre offensichtlich eine Verteilung mit einer linken oder negativen Schiefe zu bevorzugen, da es eine höhere Wahrscheinlichkeit für höhere Erträge (Bild 2.22) im Vergleich zu einer höheren Wahrscheinlichkeit für niedrigere Erträge (Bild 2.23) gibt. In einer schiefen Verteilung liefert der Medianwert deshalb eine bessere Messung der Erträge, weil die Medianwerte sowohl für Bild 2.22 als auch Bild 2.23 identisch sind, die Risiken sind identisch und daher ist ein Projekt mit einer negativ schiefen Verteilung der Nettogewinne eine bessere Wahl. Die Nichtberücksichtigung der Verteilungsschiefe eines Projekts könnte bedeuten, dass man das falsche Projekt auswählt (z.B., zwei Projekte könnten identische erste und zweite Momente aufweisen, das heißt, beide haben identische Ertrags- und Risikoprofile, aber ihre Verteilungsschiefen könnten sehr unterschiedlich sein).

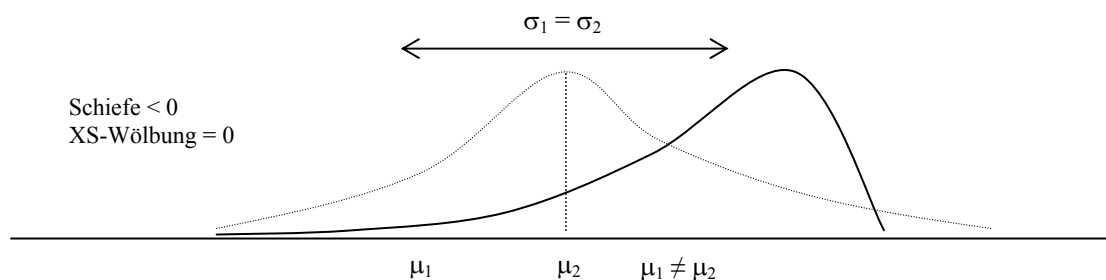


Bild 2.22—Drittes Moment (linke Schiefe)

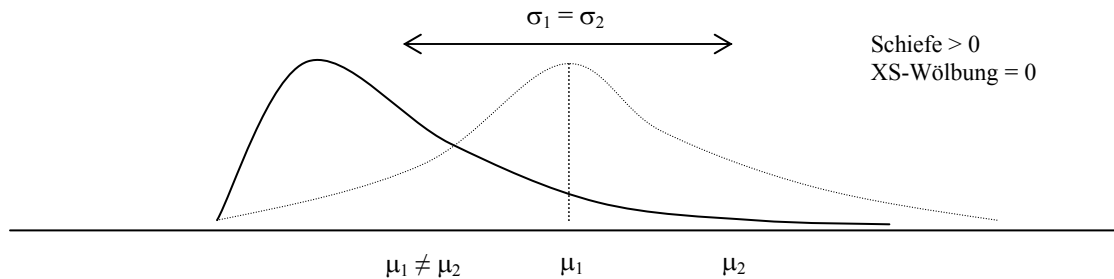


Bild 2.23—Drittes Moment (rechte Schiefe)

Vermessen der katastrophalen Schwanzereignissen in einer Verteilung — das vierte Moment

Das vierte Moment oder Wölbung misst die Spitzigkeit einer Verteilung. Bild 2.24 stellt diesen Effekt dar. Der Hintergrund (gekennzeichnet durch die punktierte Linie) ist eine Normalverteilung mit einer Wölbung von 3.0, oder einer Exzesswölbung (XS-Wölbung) von 0.0. Die Ergebnisse von Risiko Simulator zeigen den XS-Wölbung Wert, mit 0 als das normale Wölbungsniveau. Dies bedeutet, dass eine negative XS-Wölbung flachere Schwänze (platykurtische Verteilungen wie die Uniformverteilung) zeigt, während positive Werte dickere Schwänze (leptokurtische Verteilungen wie die Student-T- oder die Log-Normalverteilungen) zeigen. Die durch die Fettlinie dargestellte Verteilung hat eine höhere Exzesswölbung, deshalb ist der Bereich unter der Kurve dicker an den Schwänzen mit geringerem Bereich im Zentralkörper. Dieser Zustand hat wesentliche Auswirkungen auf die Risikoanalyse was die zwei Verteilungen in Bild 2.24 betrifft: Die ersten drei Momente (Mittelwert, Standardabweichung und Schiefe) können identisch sein, aber das vierte Moment (Wölbung) ist anders. Dieser Zustand bedeutet, dass obwohl die Erträge und die Risiken identisch sind, sind die Wahrscheinlichkeiten des Auftretens von extremen und katastrophalen Ereignissen (potentielle große Verluste oder große Gewinne) höher für eine Verteilung mit hoher Wölbung (z.B., Aktienmarktrenditen sind leptokurtisch oder haben eine hohe Wölbung). Das Ignorieren der Wölbung eines Projekts könnte nachteilig sein. Normalerweise deutet ein höherer Exzesswölbungswert darauf hin, dass die Nachteilrisiken höher sind (z.B., der Risikopotential-Wert (VAR) eines Projekts könnte signifikant sein).

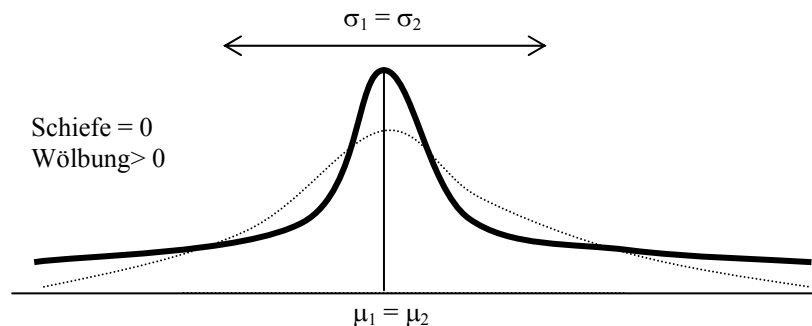


Bild 2.24—Viertes Moment

Die Funktionen der Momente

Haben Sie sich mal gefragt, warum diese Risikostatistiken, Momente genannt sind? In der mathematischen Fachsprache, bedeutet Moment, etwas erhöht zu der n-ten Potenz. In anderen Worten, das dritte Moment impliziert, dass in einer Gleichung, ist drei die am wahrscheinlichsten höchste Potenz. Die folgenden Gleichungen erläutern die mathematischen Funktionen und Anwendungen einiger Momente einer Stichprobenstatistik. Zum Beispiel, bitte bemerken Sie, dass die höchste Potenz für den Durchschnitt des ersten Moments eins, für die Standardabweichung des zweiten Moments zwei, für die Schiefe des dritten Moments drei und für die des vierten Moments vier ist.

Erstes Moment: Arithmetischer Durchschnitt oder einfacher Mittelwert (Stichprobe)

$$\bar{x} = \frac{\sum_{i=1}^n x_i}{n} \quad \text{Die äquivalente Excelfunktion ist AVERAGE}$$

Zweites Moment: Standardabweichung (Stichprobe)

$$s = \sqrt{\frac{\sum_{i=1}^n (x_i - \bar{x})^2}{n-1}} \quad \text{Die äquivalente Excelfunktion ist STDEV für eine Stichproben-}$$

standardabweichung

Die äquivalente Excelfunktion ist STDEVP für eine Bevölkerungs- standardabweichung

Drittes Moment: Schiefe

$$skew = \frac{n}{(n-1)(n-2)} \sum_{i=1}^n \frac{(x_i - \bar{x})^3}{s} \quad \text{Die äquivalente Excelfunktion ist SKEW}$$

Viertes Moment: Wölbung

$$kurtosis = \frac{n(n+1)}{(n-1)(n-2)(n-3)} \sum_{i=1}^n \frac{(x_i - \bar{x})^4}{s} - \frac{3(n-1)^2}{(n-2)(n-3)}$$

Die äquivalente Excelfunktion ist KURT

Wahrscheinlichkeitsverteilungen für Monte-Carlo-Simulationen begreifen

Dieser Abschnitt demonstriert die Stärke von Monte-Carlo-Simulationen. Um jedoch mit der Simulation zu beginnen, muss man erst das Konzept der Wahrscheinlichkeitsverteilungen begreifen. Als Anfangsschritt zum begreifen der Wahrscheinlichkeit, betrachten Sie dieses Beispiel: Sie möchten die Verteilung von nicht befreiten Gehältern innerhalb einer Abteilung eines großen Unternehmens anschauen. Als Erstes sammeln Sie die Rohdaten — in diesem Fall, die Gehälter jedes nicht befreiten Arbeitnehmers in der Abteilung. Zweitens ordnen Sie die Daten in ein sinnvolles Format ein und kartieren Sie diese Daten als eine Frequenzverteilung in einem Diagramm. Um die Frequenzverteilung zu kreieren, teilen Sie erst die Gehälter in Gruppenintervalle und dann listen Sie diese Intervalle in der Horizontalachse des Diagramms auf. Danach listen Sie die Anzahl oder Frequenz der Arbeitnehmer in jedem Intervall in der Vertikalachse des Diagramms auf. Jetzt können Sie mühelos die Verteilung von nicht befreiten Gehältern innerhalb der Abteilung anschauen. Ein Blick auf das im Bild 2.25 dargestellte Diagramm legt offen, dass die meisten Arbeitnehmer (zirka 60 aus einer Gesamtzahl von 180) zwischen \$7.00 und \$9.00 pro Stunde verdienen.

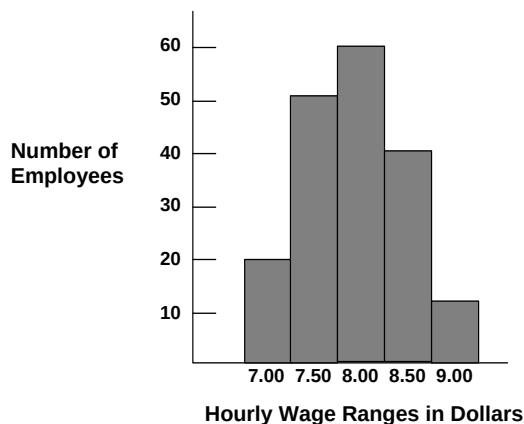


Bild 2.25—Frequenzhistogramm I

Sie können diese Daten als eine Wahrscheinlichkeitsverteilung graphisch darstellen (kartieren). Eine Wahrscheinlichkeitsverteilung zeigt die Anzahl der Arbeitnehmer in jedem Intervall als einen Bruch der Gesamtanzahl der Arbeitnehmer. Um eine Wahrscheinlichkeitsverteilung zu kreieren, teilen Sie erst die Anzahl der Arbeitnehmer in jedem Intervall durch die Gesamtanzahl der Arbeitnehmer und dann listen Sie die Ergebnisse in der Vertikalachse des Diagramms auf.

Das Diagramm im Bild 2.26 zeigt Ihnen die Anzahl der Arbeitnehmer in jeder Gehaltsgruppe als einen Bruch aller Arbeitnehmer; Sie können die Likelihood oder Wahrscheinlichkeit schätzen, dass ein zufällig aus der Gesamtgruppe ausgewählter Arbeitnehmer ein Gehalt innerhalb eines bestimmten Intervalls verdient. Zum Beispiel, angenommen dass die gleichen Bedingungen wie im Moment der Stichprobenentnahme vorhanden sind, ist die Wahrscheinlichkeit 0.33 (eine aus drei Chancen), dass ein zufällig aus der Gesamtgruppe ausgewählter Arbeitnehmer zwischen \$8.00 und \$8.50 pro Stunde verdient.

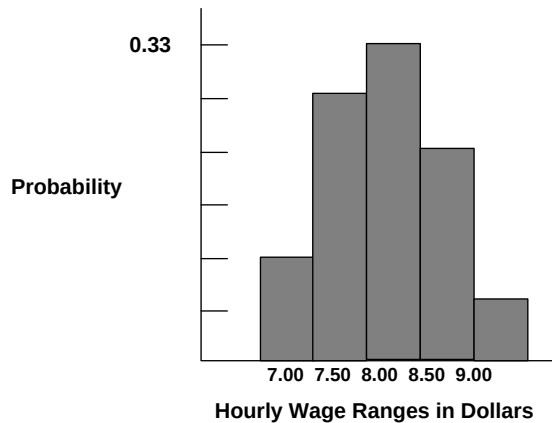


Bild 2.26—Frequenzhistogramm II

Wahrscheinlichkeitsverteilungen sind entweder diskret oder kontinuierlich. *Diskrete Wahrscheinlichkeitsverteilungen* beschreiben eindeutige Wert, normalerweise Ganzzahlen, mit keinen Zwischenwerten und werden als eine Reihe von vertikalen Balken angezeigt. Eine diskrete Verteilung, zum Beispiel, könnte die Anzahl von Köpfen in vier Münzwürfen als 0, 1, 2, 3, oder 4 beschreiben. *Kontinuierliche Verteilungen* sind eigentlich mathematische Abstraktionen, weil sie die Existenz von jedem möglichen Zwischenwert zwischen zwei Zahlen annehmen. Das heißt, eine kontinuierliche Verteilung nimmt an, dass es eine unendliche Anzahl von Werten zwischen zwei beliebigen Punkten in einer Verteilung gibt. In vielen Situationen, jedoch, können Sie tatsächlich eine kontinuierliche Verteilung verwenden, um eine diskrete Verteilung zu approximieren, obwohl das kontinuierliche Modell die Lage nicht unbedingt exakt beschreibt.

Auswahl der richtigen Wahrscheinlichkeitsverteilung

Das Plotten der Daten ist eine Hilfslinie bei der Auswahl einer Wahrscheinlichkeitsverteilung. Die folgenden Schritte erläutern ein weiteres Verfahren zur Auswahl der Wahrscheinlichkeitsverteilungen, welche die ungewissen Variablen in Ihren Tabellenblättern am besten beschreiben.

Um die richtige Wahrscheinlichkeitsverteilung auszuwählen, verwenden Sie die folgenden Schritte:

- Betrachten Sie die in Frage kommende Variable. Listen Sie alles was Sie über die Bedingungen dieser Variablen wissen auf. Vielleicht können Sie hilfreiche Informationen über die ungewisse Variable von historischen Daten sammeln. Wenn historische Daten nicht verfügbar sind, verwenden Sie Ihr eigenes Urteilsvermögen: Auf Erfahrung beruhend, listen Sie alles was Sie über die ungewisse Variable wissen auf.
- Überprüfen Sie die Beschreibungen der Wahrscheinlichkeitsverteilungen.
- Wählen Sie die Verteilung, welche diese Variable charakterisiert. Eine Verteilung charakterisiert eine Variable, wenn die Bedingungen dieser Verteilung mit denen der Variablen übereinstimmen.

Monte-Carlo-Simulation

Die Monte-Carlo-Simulation in ihrer einfachsten Form ist ein Zufallszahlengenerator, der für die Vorausberechnung, Schätzung und Risikoanalyse nützlich ist. Eine Simulation berechnet zahlreiche Szenarien eines Modells, indem sie wiederholt Werte aus einer benutzervordefinierten *Wahrscheinlichkeitsverteilung* für die ungewissen Variablen auswählt und diese Werte für das Modell verwendet. Da alle diese Szenarien dazugehörige Ergebnisse in einen Modell produzieren, kann jedes Szenario eine *Vorausberechnung* haben. Vorausberechnungen sind Ereignisse (normalerweise mit Formeln oder Funktionen), die Sie als wichtige Outputs des Modells definieren. Dies sind normalerweise Ereignisse wie Endsummen, Nettogewinn oder Bruttoausgaben.

Zur Vereinfachung, denken Sie an die Methode der Monte-Carlo-Simulation wie das wiederholte Herausnehmen mit Zurücklegung von Golfbällen aus einem großen Korb. Die Größe und Form des Korbs hängt von den *Hypothesen* der Verteilung ab (z.B., eine Normalverteilung mit einem Mittelwert von 100 und einer Standardabweichung von 10, gegen eine Uniformverteilung oder eine Dreiecksverteilung) wobei einige Körbe tiefer oder symmetrischer als andere sind, was bedeutet, dass einige Bälle häufiger als andere herausgenommen werden. Die Anzahl der wiederholt herausgenommenen Bälle hängt von der Anzahl der simulierten *Probeversuche* ab. Für ein großes Modell mit mehrfachen verwandten Hypothesen, stellen Sie sich das große Modell wie einen sehr großen Korb vor, in dem viele Minikörbe stecken. Jeder Minikorb hat seine eigene Menge von herum springenden Golfbällen. Gelegentlich stehen diese Minikörbe in Verbindung zueinander (wenn es eine *Korrelation* zwischen den Variablen gibt) und die Golfbälle springen im Zusammenhang miteinander herum, während andere unabhängig voneinander herumspringen. Die Bälle, die jedes Mal aus diesen Interaktionen innerhalb des Modells (der große zentrale Korb) herausgenommen werden, werden tabelliert und aufgezeichnet, was ein *Vorausberechnungsergebnis* der Simulation liefert.

Mit einer Monte-Carlo-Simulation generiert Risiko Simulator völlig voneinander unabhängige Zufallswerte für die Wahrscheinlichkeitsverteilung jeder Hypothese. In anderen Worten, der für einen Probeversuch ausgewählte Zufallswert hat keinen Effekt auf den als nächsten generierten Zufallswert. Verwenden Sie das Monte-Carlo-Stichprobenverfahren, wenn Sie reale „was wäre Szenarien“ für ihr Tabellenblattmodell simulieren möchten.

Diskrete Verteilungen

Es folgt eine detaillierte Auflistung der verschiedenen Typen von Wahrscheinlichkeitsverteilungen, die man in Monte-Carlo-Simulationen verwenden kann. Diese Auflistung ist für den Benutzer zum Nachschlagen in diesem Anhang beigelegt.

Bernoulli oder Ja/Nein Verteilung

Die Bernoulli-Verteilung ist eine diskrete Verteilung mit zwei Ergebnissen (z.B., Kopf oder Zahl, Erfolg oder Misserfolg, 0 oder 1). Die Bernoulli-Verteilung ist die Binomialverteilung mit einem Probeversuch und kann verwendet werden, um Ja/Nein oder Erfolg/Misserfolg Verhältnisse zu simulieren. Diese Verteilung ist der grundlegende Baustein von anderen komplexeren Verteilungen. Zum Beispiel:

- Binomialverteilung : Bernoulli-Verteilung mit einer höheren Anzahl von n Gesamtprobeversuchen und berechnet die Wahrscheinlichkeit von x Erfolgen innerhalb dieser Gesamtzahl von Probeversuchen.
- Geometrische Verteilung : Bernoulli-Verteilung mit einer höheren Anzahl von Gesamtprobeversuchen und berechnet die Anzahl der erforderlichen Misserfolge, bevor der erste Erfolg vorkommt.
- Negative Binomialverteilung: Bernoulli-Verteilung mit einer höheren Anzahl von Gesamtprobeversuchen berechnet die Anzahl der Misserfolge, bevor der n -te Erfolg vorkommt.

Die mathematischen Konstrukte für die Bernoulli-Verteilung sind wie folgt:

$$P(n) = \begin{cases} 1-p & \text{für } x=0 \\ p & \text{für } x=1 \end{cases}$$

oder

$$P(n) = p^x (1-p)^{1-x}$$

Mittelwert = p

$$\text{Standardabweichung} = \sqrt{p(1-p)}$$

$$\text{Schiefe} = \frac{1-2p}{\sqrt{p(1-p)}}$$

$$\text{Exzesswölbung} = \frac{6p^2 - 6p + 1}{p(1-p)}$$

Die Erfolgswahrscheinlichkeit (p) ist der einzige Verteilungsparameter. Es ist auch wichtig zu bemerken, dass es nur einen Probeversuch in der Bernoulli-Verteilung gibt und der resultierende simulierte Wert entweder 0 oder 1 ist.

Inputanforderungen:

Erfolgswahrscheinlichkeit > 0 und < 1 (das heißt, $0.0001 \leq p \leq 0.9999$)

Binomialverteilung

Die Binomialverteilung beschreibt die Anzahl der Male, in denen ein bestimmtes Ereignis in einer vordefinierten Anzahl von Probeversuchen stattfindet, sowie die Anzahl der Köpfe bei 10 Münzwürfen oder die Anzahl der mangelhaften Artikel aus einer Auswahl von 50 Artikeln.

Bedingungen

Die drei unterliegenden Bedingungen einer Binomialverteilung sind:

- Für jeden Probeversuch gibt es nur zwei sich gegenseitig ausschließende Ergebnisse.
- Die Probeversuche sind unabhängig—was sich während des ersten Probeversuchs ereignet hat keinen Einfluss auf den nächsten Probeversuch.
- Die Wahrscheinlichkeit des Eintritts eines Ereignis bleibt gleich von Probeversuch zu Probeversuch.

Die mathematischen Konstrukte für die Binomialverteilung sind wie folgt:

$$P(x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{(n-x)} \quad \text{für } n > 0; x = 0, 1, 2, \dots, n; \text{ und } 0 < p < 1$$

Mittelwert = np

Standardabweichung = $\sqrt{np(1-p)}$

$$\text{Schiefe} = \frac{1-2p}{\sqrt{np(1-p)}}$$

$$\text{Exzesswölbung} = \frac{6p^2 - 6p + 1}{np(1-p)}$$

Die Erfolgswahrscheinlichkeit (p) und die ganzzahlige Anzahl der Gesamtprobeversuche (n) sind die Verteilungsparameters. Die Anzahl der erfolgreichen Probeversuche wird als x bezeichnet. Es ist wichtig zu bemerken, dass die Erfolgswahrscheinlichkeit (p) von 0 oder 1 eine geringfügige Bedingung ist und keine Simulationen benötigt, und ist deshalb nicht in der Software erlaubt.

Inputanforderungen:

Erfolgswahrscheinlichkeit > 0 und < 1 (das heißt, $0.0001 \leq p \leq 0.9999$)

Anzahl der Probeversuche ≥ 1 oder positive Ganzzahlen und ≤ 1000 (für größere Probeversuche verwenden Sie die Normalverteilung mit dem relevanten berechneten binomialen Mittelwert und die relevante berechnete binomiale Standardabweichung als die Parameter der Normalverteilung).

Diskrete Uniform

Die diskrete Uniformverteilung ist auch als die Verteilung mit *wahrscheinlich gleicher Ausgänge* bekannt, wobei wenn die Verteilung einen Satz von N Elementen hat, dann hat jedes Element die gleiche Wahrscheinlichkeit. Diese Verteilung ist mit der Uniformverteilung verwandt, aber ihre Elemente sind diskret und nicht kontinuierlich.

Die mathematischen Konstrukte für die Binomialverteilung sind wie folgt:

$$P(x) = \frac{1}{N}$$

$$\text{Mittelwert} = \frac{N+1}{2} \text{ rangierter Wert}$$

$$\text{Standardabweichung} = \sqrt{\frac{(N-1)(N+1)}{12}} \text{ rangierter Wert}$$

Schiefe = 0 (das heißt, die Verteilung ist perfekt symmetrisch)

$$\text{Exzesswölbung} = \frac{-6(N^2+1)}{5(N-1)(N+1)} \text{ rangierter Wert}$$

Inputanforderungen:

Minimum < Maximum und müssen Ganzzahlen sein (negative Ganzzahlen und Null sind erlaubt)

Geometrische Verteilung

Die geometrische Verteilung beschreibt die Anzahl der Probeversuche, bis zum ersten erfolgreichen Vorkommnis, sowie die Anzahl der Male, die Sie einen Roulettekessel drehen müssen, bis Sie gewinnen

Bedingungen

Die drei unterliegenden Bedingungen der geometrischen Verteilung sind:

- Die Anzahl der Probeversuche ist nicht festgelegt.
- Die Probeversuche dauern bis zum ersten Erfolg an.
- Die Erfolgswahrscheinlichkeit ist die gleiche von Probeversuch zu Probeversuch.

Die mathematischen Konstrukte für die geometrische Verteilung sind wie folgt:

$$P(x) = p(1-p)^{x-1} \text{ für } 0 < p < 1 \text{ und } x = 1, 2, \dots, n$$

$$\text{Mittelwert} = \frac{1}{p} - 1$$

$$\text{Standardabweichung} = \sqrt{\frac{1-p}{p^2}}$$

$$\text{Schiefe} = \frac{2-p}{\sqrt{1-p}}$$

$$\text{Exzesswölbung} = \frac{p^2 - 6p + 6}{1-p}$$

Die Erfolgswahrscheinlichkeit (p) ist der einzige Verteilungsparameter. Die Anzahl der

erfolgreichen simulierten Probeversuche wird als x bezeichnet und kann nur aus positiven Ganzzahlen bestehen.

Inputanforderungen:

Erfolgswahrscheinlichkeit > 0 und < 1 (das heißt, $0.0001 \leq p \leq 0.9999$). Es ist wichtig zu bemerken, dass die Erfolgswahrscheinlichkeit (p) von 0 oder 1 eine geringfügige Bedingung ist und keine Simulationen benötigt, und ist deshalb nicht in der Software erlaubt.

Hypergeometrische Verteilung

Die hypergeometrische Verteilung ähnelt der Binomialverteilung, da beide die Anzahl der Male beschreiben, in denen sich ein bestimmtes Ereignis in einer festgelegten Anzahl von Probenversuchen ereignet. Der Unterschied ist, dass die Probeversuche der Binomialverteilung unabhängig sind, während die Probeversuche der hypergeometrischen Verteilung die Wahrscheinlichkeit für jeden folgenden Probeversuch ändern und werden als "Probeversuche ohne Zurücklegung" bezeichnet. Nehmen wir zum Beispiel an, dass ein Karton von hergestellten Teilen bekanntermaßen einige mangelhafte Teile enthält. Sie wählen ein Teil aus dem Karton, stellen fest, dass es defekt ist und entfernen das Teil vom Karton. Wenn Sie ein weiteres Teil aus dem Karton auswählen, ist die Wahrscheinlichkeit, dass es defekt ist etwas geringer als für das erste Teil, weil sie ein mangelhaftes Teil entfernt hatten. Hätten Sie das mangelhafte Teil zurückgelegt, wären die Wahrscheinlichkeiten gleich geblieben und der Prozess hätte die Bedingungen für eine Binomialverteilung befriedigt.

Bedingungen

Die unterliegenden Bedingungen der hypergeometrischen Verteilung sind:

- Die Gesamtzahl der Gegenstände oder Elemente (die Bevölkerungsgröße) ist eine feste Nummer, eine endliche Bevölkerung. Die Bevölkerungsgröße muss weniger als oder gleich 1,750 sein.
- Die Stichprobengröße (die Anzahl der Probeversuche) repräsentiert einen Teil der Bevölkerung.
- Die bekannte anfängliche Erfolgswahrscheinlichkeit in der Bevölkerung ändert sich nach jedem Probeversuch.

Die mathematischen Konstrukte für die hypergeometrische Verteilung sind wie folgt:

$$P(x) = \frac{\frac{(N_x)!}{x!(N_x - x)!} \frac{(N - N_x)!}{(n - x)!(N - N_x - n + x)!}}{\frac{N!}{n!(N - n)!}} \quad \text{für } x = \text{Max}(n - (N - N_x), 0), \dots, \text{Min}(n, N_x)$$

$$\text{Mittelwert} = \frac{N_x n}{N}$$

$$\text{Standardabweichung} = \sqrt{\frac{(N - N_x) N_x n (N - n)}{N^2 (N - 1)}}$$

$$\text{Schiefe} = \sqrt{\frac{N-1}{(N-N_x)N_x n(N-n)}}$$

Exzesswölbung = komplexe Funktion

Die Verteilungsparameter sind die Anzahl der Elemente in der Bevölkerung oder die Bevölkerungsgröße (N), die abgetasteten Probeversuche oder die Stichprobengröße (n) und die Anzahl der Elemente in der Bevölkerung, welche die Erfolgscharakteristik oder Bevölkerungserfolge (N_x) besitzen. Die Anzahl der erfolgreichen Probeversuche wird als x bezeichnet.

Inputanforderungen:

Bevölkerungsgröße ≥ 2 und Ganzzahl

Stichprobengröße > 0 und Ganzzahl

Bevölkerungserfolge > 0 und Ganzzahl

Bevölkerungsgröße $>$ Bevölkerungserfolge

Stichprobengröße $<$ Bevölkerungserfolge

Bevölkerungsgröße < 1750

Negative Binomialverteilung

Die negative Binomialverteilung ist nützlich bei der Modellierung der Verteilung der zusätzlichen Probeversuche, die über der Anzahl der erforderlichen Erfolgsvorkommnisse (R) erfordert wird. Zum Beispiel, um eine gesamte Menge von 10 erfolgreichen Verkaufschancen, wie viele extra Verkaufbesuche müsste man über die 10 Besuche machen, gegeben eine Erfolgswahrscheinlichkeit bei jedem Besuch? Die x-Achse zeigt die Anzahl der erforderlichen zusätzlichen Besuche oder die Anzahl der erfolglosen Besuche an. Die Anzahl der Probeversuche ist nicht fest, die Probeversuche dauern bis zum R-ten Erfolg an und die Erfolgswahrscheinlichkeit ist die gleiche von Probeversuch zu Probeversuch. Die Erfolgswahrscheinlichkeit (p) und die Anzahl der erforderlichen Erfolge (R) sind die Verteilungsparameter. Es handelt sich eigentlich um eine *Superverteilung* der geometrischen und Binomialverteilungen. Diese Verteilung zeigt die Wahrscheinlichkeiten von jeder Anzahl der Probeversuche, die mehr als R sind, um den erforderlichen Erfolg R zu liefern.

Bedingungen

Die drei unterliegenden Bedingungen der negativen Binomialverteilung sind:

- Die Anzahl der Probeversuche ist nicht festgelegt.
- Die Probeversuche dauern bis zum r -ten Erfolg an.
- Die Erfolgswahrscheinlichkeit ist die gleiche von Probeversuch zu Probeversuch.

Die mathematischen Konstrukte für die negative Binomialverteilung sind wie folgt:

$$P(x) = \frac{(x+r-1)!}{(r-1)!x!} p^r (1-p)^x \quad \text{für } x = r, r+1, \dots; \text{ und } 0 < p < 1$$

$$\text{Mittelwert} = \frac{r(1-p)}{p}$$

$$\text{Standardabweichung} = \sqrt{\frac{r(1-p)}{p^2}}$$

$$\text{Schiefe} = \frac{2-p}{\sqrt{r(1-p)}}$$

$$\text{Exzesswölbung} = \frac{p^2 - 6p + 6}{r(1-p)}$$

Erfolgswahrscheinlichkeit (p) und erforderliche Erfolge (R) sind die Verteilungsparameter.

Inputanforderungen:

Erforderliche Erfolge müssen positive Ganzzahlen > 0 und < 8000 .

Erfolgswahrscheinlichkeit > 0 und < 1 (das heißt, $0.0001 \leq p \leq 0.9999$). Es ist wichtig zu bemerken, dass die Erfolgswahrscheinlichkeit (p) von 0 oder 1 eine geringfügige Bedingung ist und keine Simulationen benötigt, und ist deshalb nicht in der Software erlaubt.

Pascal-Verteilung

Die Pascal-Verteilung ist hilfreich bei der Modellierung der Verteilung der Gesamtanzahl der benötigten Probeversuche, um die Anzahl der erforderlichen erfolgreichen Vorkommnisse zu erreichen. Zum Beispiel, um 10 Verkaufsgelegenheiten erfolgreich abzuschließen, wie viele Verkaufsanrufe insgesamt müssten man führen, unter Berücksichtigung einer bestimmten Erfolgswahrscheinlichkeit bei jedem Anruf? Die x-Achse zeigt die Gesamtanzahl der erforderlichen Anrufe, welche sowohl erfolgreiche als auch erfolglose Anrufe einschließt. Die Anzahl der Probeversuche ist nicht festgelegt, die Probeversuche dauern bis zum R-ten Erfolg an und die Erfolgswahrscheinlichkeit bleibt gleich von Probeversuch zu Probeversuch. Die Pascal-Verteilung ist mit der negativen Binomialverteilung verwandt. Die negative Binomialverteilung kalkuliert die Anzahl der Ereignisse, die über die Anzahl der gewünschten Erfolge benötigt wird, unter Berücksichtigung einer bestimmten Wahrscheinlichkeit (in anderen Worten, die Gesamtausfälle). Die Pascal-Verteilung, hingegen, kalkuliert die Gesamtanzahl der erforderlichen Ereignisse (in anderen Worten, die Summe der Ausfälle und der Erfolge), um die gewünschten Erfolge zu erreichen, unter Berücksichtigung einer bestimmten Wahrscheinlichkeit. Erforderliche Erfolge und Wahrscheinlichkeit sind die Verteilungsparameter.

Die mathematischen Konstrukte für die Pascal sind wie folgt:

$$f(x) = \begin{cases} \frac{(x-1)!}{(x-s)!(s-1)!} p^s (1-p)^{x-s} & \text{for all } x \geq s \\ 0 & \text{otherwise} \end{cases}$$

$$F(x) = \begin{cases} \sum_{x=1}^k \frac{(x-1)!}{(x-s)!(s-1)!} p^s (1-p)^{x-s} & \text{for all } x \geq s \\ 0 & \text{otherwise} \end{cases}$$

$$\text{Mittelwert} = \frac{s}{p}$$

$$\text{Standardabweichung} = \sqrt{s(1-p)p^2}$$

$$\text{Schiefe} = \frac{2-p}{\sqrt{r(1-p)}}$$

$$\text{Exzesswölbung} = \frac{p^2 - 6p + 6}{r(1-p)}$$

Inputvoraussetzungen:

Erforderliche Erfolge > 0 und muss eine Ganzzahl sein

$0 \leq \text{Wahrscheinlichkeit} \leq 1$

Poisson-Verteilung

Die Poisson-Verteilung beschreibt die Anzahl der Male, in denen ein Ereignis in einem gegebenen Intervall stattfindet, sowie die Anzahl der Telefonanrufe pro Minute oder die Anzahl der Fehler pro Seite in einem Dokument

Bedingung

Die drei unterliegenden Bedingungen der Poisson-Verteilung sind:

- Die Anzahl der möglichen Vorkommen in einem Intervall ist unbegrenzt.
- Die Vorkommen sind unabhängig. Die Anzahl der Vorkommen in einem Intervall beeinflussen nicht die Anzahl der Vorkommen in anderen Intervallen.
- Die Durchschnittsanzahl der Vorkommen muss gleich bleiben von Intervall zu Intervall.

Die mathematischen Konstrukte für die Poisson sind wie folgt:

$$P(x) = \frac{e^{-\lambda} \lambda^x}{x!} \quad \text{für } x \text{ und } \lambda > 0$$

$$\text{Mittelwert} = \lambda$$

$$\text{Standardabweichung} = \sqrt{\lambda}$$

$$\text{Schiefe} = \frac{1}{\sqrt{\lambda}}$$

$$\text{Exzesswölbung} = \frac{1}{\lambda}$$

Die Rate oder Lambda (λ) ist der einzige Verteilungsparameter.

Inputanforderungen:

Rate > 0 und ≤ 1000 (das heißt, $0.0001 \leq \text{Rate} \leq 1000$)

Kontinuierliche Verteilungen

Arcussinus-Verteilung

Die Arcussinus-Verteilung ist u-förmig und ist ein spezieller Fall der Beta-Verteilung, wenn sowohl Form und Skala gleich 0.5 sind. Werte, die nahe beim Minimum und beim Maximum liegen, haben hohe Auftrittswahrscheinlichkeiten, während Werte die sich zwischen diesen beiden Extrema befinden, nur sehr kleine Auftrittswahrscheinlichkeiten haben. Minimum und Maximum sind die Verteilungsparameter.

Die mathematischen Konstrukte für die arcsine sind wie folgt:

$$f(x) = \begin{cases} \frac{1}{\pi\sqrt{x(1-x)}} & \text{for } 0 \leq x \leq 1 \\ 0 & \text{otherwise} \end{cases}$$

$$F(x) = \begin{cases} 0 & x < 0 \\ \frac{2}{\pi} \sin^{-1}(\sqrt{x}) & \text{for } 0 \leq x \leq 1 \\ 1 & x > 1 \end{cases}$$

$$\text{Mittelwert} = \frac{\text{Min} + \text{Max}}{2}$$

$$\text{Standardabweichung} = \sqrt{\frac{(\text{Max} - \text{Min})^2}{8}}$$

Schiefe = 0

Exzesswölbung = 1.5

Inputvoraussetzungen:

Minimum < Maximum

Betaverteilung

Die Betaverteilung ist sehr flexibel und wird üblicherweise verwendet, um die Streuung über einen festgestellten Bereich zu repräsentieren. Eine der wichtigsten Anwendungen der Betaverteilung ist ihre Verwendung als eine konjugierte Verteilung für den Parameter einer Bernoulli-Verteilung. In dieser Anwendung wird die Betaverteilung verwendet, um die Ungewissheit in der Wahrscheinlichkeit des Auftretens eines Ereignisses zu repräsentieren. Sie wird auch verwendet, um empirische Daten zu beschreiben und um das Zufallsverhalten von Prozentsätzen und Brüchen vorherzusagen, da sich der Bereich der Ergebnisse typisch zwischen 0

und 1 befindet. Der Wert der Betaverteilung liegt in der Formenvielfalt, die sie annehmen kann, wenn man die zwei Parameter, Alpha und Beta, variiert. Wenn die Parameter gleich sind, ist die Verteilung symmetrisch. Wenn einer der beiden Parameter 1 ist und der andere Parameter größer als 1 ist, ist die Verteilung J-förmig. Wenn Alpha kleiner als Beta ist, wird die Verteilung als positiv asymmetrisch bezeichnet (die meisten Werte liegen in der Nähe des Minimalwerts). Wenn Alpha größer als Beta ist, ist die Verteilung negativ asymmetrisch (die meisten Werte liegen in der Nähe des Maximalwerts).

Die mathematischen Konstrukte für die Betaverteilung sind wie folgt:

$$f(x) = \frac{(x)^{(\alpha-1)} (1-x)^{(\beta-1)}}{\left[\frac{\Gamma(\alpha)\Gamma(\beta)}{\Gamma(\alpha+\beta)} \right]} \quad \text{für } \alpha > 0; \beta > 0; x > 0$$

$$\text{Mittelwert} = \frac{\alpha}{\alpha + \beta}$$

$$\text{Standardabweichung} = \sqrt{\frac{\alpha\beta}{(\alpha + \beta)^2 (1 + \alpha + \beta)}}$$

$$\text{Schiefe} = \frac{2(\beta - \alpha)\sqrt{1 + \alpha + \beta}}{(2 + \alpha + \beta)\sqrt{\alpha\beta}}$$

$$\text{Exzesswölbung} = \frac{3(\alpha + \beta + 1)[\alpha\beta(\alpha + \beta - 6) + 2(\alpha + \beta)^2]}{\alpha\beta(\alpha + \beta + 2)(\alpha + \beta + 3)} - 3$$

Alpha (α) und Beta (β) sind die zwei Verteilungsformparameter und Γ ist die Gammafunktion.

Bedingungen

Die zwei unterliegenden Bedingungen der Betaverteilung sind:

- Die Ungewisse Variable ist ein Zufallswert zwischen 0 und einem positiven Wert.
- Man kann die Form der Verteilung unter Verwendung von zwei positiven Werten spezifizieren.

Inputanforderungen:

Alpha und Beta sind beide > 0 und können einen beliebigen positiven Wert haben

Multiplikative Beta versetzte Verteilung

Die Betaverteilung ist sehr flexibel und wird häufig verwendet, um die Veränderlichkeit über einem festgestellten Intervall zu repräsentieren. Sie wird verwendet, um empirische Daten zu beschreiben und das Zufallsverhalten von Prozentualen und Brüchen vorherzusagen, da das Ausgangsintervall normalerweise zwischen 0 und 1 liegt. Der Wert der Betaverteilung liegt in der großen Vielfalt der Formen, die sie annehmen kann, wenn man die zwei Parameter, Alpha und Beta, variiert. Die Beta versetzte Verteilung wird erhalten, indem man die Betaverteilung mit einem Faktor multipliziert und die Ergebnisse um einen bestimmten Lageparameter versetzt, um es dem Intervall der Ausgänge zu ermöglichen, sich über die natürlichen Grenzen von 0 und 1 zu

erweitern mit einem anderen Ausgangspunkt als 0. Alpha, Beta, Lage und Faktor sind die Inputparameter.

Inputvoraussetzungen:

Alpha > 0

Beta > 0

Lage kann jeder beliebiger positiver oder negativer Wert, einschließlich Null, sein

Faktor > 0

Cauchy- oder Lorentz- oder Breit-Wigner-Verteilung

Die Cauchy-Verteilung, auch als die Lorentz- oder Breit-Wigner-Verteilung bezeichnet, ist eine kontinuierliche Verteilung, die das Resonanzverhalten beschreibt. Sie beschreibt auch die Verteilung der horizontalen Entfernungen bei denen ein Liniensegment, das mit einem Zufallswinkel geneigt ist, die x-Achse schneidet.

Die mathematischen Konstrukte für die Cauchy- oder Lorentz-Verteilung sind wie folgt:

$$f(x) = \frac{1}{\pi} \frac{\gamma / 2}{(x - m)^2 + \gamma^2 / 4}$$

Die Cauchy-Verteilung ist ein Spezialfall: Sie hat keine theoretische Momente (Mittelwert, Standardabweichung, Schiefe und Wölbung, da sie alle undefiniert sind.

Die Modalwertposition (α) und Skala (β) sind die einzigen zwei Parameter in dieser Verteilung. Der Positionsparameter spezifiziert die Spitze oder den Modalwert der Verteilung, während der Skalenparameter die Halbspannweite beim Halbmaximum der Verteilung angibt. Ferner sind der Mittelwert und die Varianz einer Cauchy- oder Lorentz-Verteilung undefiniert.

Außerdem ist die Cauchy-Verteilung die Studentsche T-Verteilung mit nur einem 1 Freiheitsgrad. Diese Verteilung wird auch folgendermaßen aufgestellt: Man nimmt das Verhältnis von zwei Standardnormalverteilungen (Normalverteilungen mit einem Mittelwert von Null und einer Varianz von 1), die unabhängig voneinander sind.

Inputanforderungen:

Position Alpha kann irgendeinen Wert haben

Skala Beta > 0 und kann jeden beliebigen positiven Wert haben

Chi-Quadrat Verteilung

Die Chi-Quadrat Verteilung ist eine Wahrscheinlichkeitsverteilung, die hauptsächlich beim Hypothesentesten verwendet wird. Sie ist mit der Gammaverteilung und der Standardnormalverteilung verwandt. Zum Beispiel, die Summe der unabhängigen Normalverteilungen werden als eine Chi-Quadrat (χ^2) mit k Freiheitsgraden verteilt:

$$Z_1^2 + Z_2^2 + \dots + Z_k^2 \stackrel{d}{\sim} \chi_k^2$$

Die mathematischen Konstrukte für die Chi-Quadrat Verteilung sind wie folgt:

$$f(x) = \frac{0.5^{-k/2}}{\Gamma(k/2)} x^{k/2-1} e^{-x/2} \text{ für alle } x > 0$$

Mittelwert = k

Standardabweichung = $\sqrt{2k}$

$$\text{Schiefe} = 2\sqrt{\frac{2}{k}}$$

$$\text{Exzesswölbung} = \frac{12}{k}$$

Γ ist die Gammafunktion. Freiheitsgrade k sind die einzigen Verteilungsparameter.

Man kann die Chi-Quadrat Verteilung auch unter Verwendung einer Gammaverteilung modellieren, indem man folgendes einstellt:

Formparameter = $\frac{k}{2}$ und Skala = $2S^2$, wobei S die Skala ist.

Inputanforderungen:

Freiheitsgrade > 1 und muss eine Ganzzahl < 300 sein

Doppelte Log-Verteilung

Die doppelte Log-Verteilung sieht wie eine Cauchy-Verteilung aus, wobei die Zentraltendenz gespitzt ist und den maximalen Wert der Wahrscheinlichkeitsdichte trägt. Er vermindert sich aber schneller, je mehr er sich vom Zentrum entfernt, was eine symmetrische Verteilung mit einer extremen Spitze zwischen den maximalen und den minimalen Werten kreiert. Minimum und Maximum sind die Verteilungsparameter.

Die mathematischen Konstrukte für die doppelte log sind wie folgt:

$$f(x) = \begin{cases} \frac{-1}{2b} \ln\left(\frac{|x-a|}{b}\right) & \text{for } \min \leq x \leq \max \\ 0 & \text{otherwise} \end{cases}$$

$$\text{where } a = \frac{\min + \max}{2} \text{ and } b = \frac{\max - \min}{2}$$

$$F(x) = \begin{cases} \frac{1}{2} - \left(\frac{|x-a|}{2b}\right) \left[1 - \ln\left(\frac{|x-a|}{b}\right)\right] & \text{for } \min \leq x \leq a \\ \frac{1}{2} + \left(\frac{|x-a|}{2b}\right) \left[1 - \ln\left(\frac{|x-a|}{b}\right)\right] & \text{for } a \leq x \leq \max \end{cases}$$

$$\text{Mittelwert} = \frac{\text{Min} + \text{Max}}{2}$$

$$\text{Standardabweichung} = \sqrt{\frac{(\text{Max} - \text{Min})^2}{36}}$$

$$\text{Schiefe} = 0$$

Inputvoraussetzungen:

Minimum < Maximum

Dreiecksverteilung

Die Dreiecksverteilung beschreibt eine Situation wo man den Minimum, den Maximum und den am wahrscheinlichsten der auftretenden Werte kennt. Zum Beispiel, Sie könnten die Anzahl der pro Woche verkauften Autos beschreiben, wenn frühere Verkäufe die minimale, maximale und normale Anzahl der verkauften Autos anzeigen.

Bedingungen

Die drei unterliegenden Bedingungen der Dreiecksverteilung sind:

- Die minimale Anzahl der Elemente ist festgelegt.
- Die maximale Anzahl der Elemente ist festgelegt.
- Die wahrscheinlichste Anzahl der Elemente fällt zwischen die Minimal- und Maximalwerte, was eine dreieckförmige Verteilung bildet. Diese zeigt, dass sich die Werte in der Nähe des Minimums und Maximums mit geringerer Wahrscheinlichkeit ereignen, als die in der Nähe des wahrscheinlichsten Wertes.

Die mathematischen Konstrukte für die Dreiecksverteilung sind wie folgt:

$$f(x) = \begin{cases} \frac{2(x - \text{Min})}{(\text{Max} - \text{Min})(\text{Likely} - \text{Min})} & \text{für } \text{Min} < x < \text{Likely} \\ \frac{2(\text{Max} - x)}{(\text{Max} - \text{Min})(\text{Max} - \text{Likely})} & \text{für } \text{Likely} < x < \text{Max} \end{cases}$$

$$\text{Mittelwert} = \frac{1}{3}(\text{Min} + \text{Likely} + \text{Max})$$

$$\text{Standardabweichung} = \sqrt{\frac{1}{18}(\text{Min}^2 + \text{Likely}^2 + \text{Max}^2 - \text{Min Max} - \text{Min Likely} - \text{Max Likely})}$$

$$\text{Schiefe} = \frac{\sqrt{2}(\text{Min} + \text{Max} - 2\text{Likely})(2\text{Min} - \text{Max} - \text{Likely})(\text{Min} - 2\text{Max} + \text{Likely})}{5(\text{Min}^2 + \text{Max}^2 + \text{Likely}^2 - \text{MinMax} - \text{MinLikely} - \text{MaxLikely})^{3/2}}$$

Exzesswölbung = -0.6 (dies gilt für alle Inputs von *Min*, *Max* und *wahrscheinlichster Wert (Likely)*)

Minimalwert (*Min*), wahrscheinlichster Wert (*Likely*) und Maximalwert (*Max*) sind die Verteilungsparameter.

Inputanforderungen:

$\text{Min} \leq \text{wahrscheinlichster} \leq \text{Max}$ und können beliebige Werte sein

Allerdings, $\text{Min} < \text{Max}$ und können beliebige Werte sein

Erlang-Verteilung

Die Erlang-Verteilung ist gleich der Gammaverteilung mit der Voraussetzung, dass Alpha, oder der Formparameter, eine positive Ganzzahl sein muss. Eine Beispielsanwendung der Erlang-Verteilung ist die Kalibrierung der Transitionsrate von Elementen durch ein System von Unterteilungen. Solche Systeme werden häufig in Biologie und Ökologie verwendet (z.B., bei der Epidemiologie: Eine Person könnte sich mit einer exponentiellen Rate vom Gesundsein zum Krankheitsträger entwickeln, und exponentiell vom Träger zur Infektionsgefahr fortschreiten). Alpha (auch als Form bekannt) und Beta (auch als Skala bekannt) sind die Verteilungsparameter.

Die mathematischen Konstrukte für die Erlang sind wie folgt:

$$f(x) = \begin{cases} \left(\frac{x}{\beta}\right)^{\alpha-1} e^{-x/\beta} & \text{for } x \geq 0 \\ \frac{\beta(\alpha-1)}{\beta(\alpha-1)} & \\ 0 & \text{otherwise} \end{cases}$$

$$F(x) = \begin{cases} 1 - e^{-x/\beta} \sum_{i=0}^{\alpha-1} \frac{(x/\beta)^i}{i!} & \text{for } x \geq 0 \\ 0 & \text{otherwise} \end{cases}$$

Mittelwert = $\alpha\beta$

Standardabweichung = $\sqrt{\alpha\beta^2}$

Schiefe = $\frac{2}{\sqrt{\alpha}}$

Exzesswölbung = $\frac{6}{\alpha} - 3$

Inputvoraussetzungen:

Alpha (Form) > 0 und muss eine Ganzzahl sein

Beta (Skala) > 0

Exponentialverteilung

Die Exponentialverteilung wird verbreitet verwendet, um Ereignisse zu beschreiben, die an Zufallszeitpunkten stattfinden, sowie der Zeitrahmen zwischen Ereignissen wie Ausfälle von elektronischen Geräten oder die Zeit zwischen Ankünften an einem Serviceschalter. Sie ist mit der Poisson-verteilung verwandt, welche die Eintrittszahl eines Ereignisses innerhalb eines gegebenen Zeitraums beschreibt. Eine wichtige Eigenschaft der Exponentialverteilung ist ihre

“Gedächtnislosigkeit”, was bedeutet, dass die zukünftige Lebensdauer eines gegebenen Objektes dieselbe Verteilung hat, unabhängig von der Zeit, in der es existiert hat. In anderen Worten, die Zeit hat keine Auswirkung auf zukünftige Ausgänge.

Die mathematischen Konstrukte für die Exponentialverteilung sind wie folgt:

$$f(x) = \lambda e^{-\lambda x} \quad \text{für } x \geq 0; \lambda > 0$$

$$\text{Mittelwert} = \frac{1}{\lambda}$$

$$\text{Standardabweichung} = \frac{1}{\lambda}$$

Schiefe = 2 (dieser Wert gilt für all Inputs der Erfolgsrate λ)

Exzesswölbung = 6 (dieser Wert gilt für all Inputs der Erfolgsrate λ)

Erfolgsrate (λ) ist der einzige Verteilungsparameter. Die Anzahl der erfolgreichen Probeversuche wird als x bezeichnet.

Bedingungen

Die unterliegende Bedingung der Exponentialverteilung ist:

- Die Exponentialverteilung beschreibt die Dauer zwischen Vorkommnisse.

Inputanforderungen:

Rate > 0

Exponentialverteilung Versetzte

Die Exponentialverteilung wird verbreitet verwendet, um Ereignisse zu beschreiben, die an Zufallszeitpunkten stattfinden, sowie der Zeitrahmen zwischen Ereignissen wie Ausfälle von elektronischen Geräten oder die Zeit zwischen Ankünften an einem Serviceschalter. Sie ist mit der Poissonverteilung verwandt, welche die Eintrittszahl eines Ereignisses innerhalb eines gegebenen Zeitraums beschreibt. Eine wichtige Eigenschaft der Exponentialverteilung ist ihre Gedächtnislosigkeit, was bedeutet, dass die zukünftige Lebensdauer eines gegebenen Objektes dieselbe Verteilung hat, unabhängig von der Zeit, in der es schon existiert hat. In anderen Worten, die Zeit hat keine Auswirkung auf zukünftige Ausgänge. Die Erfolgsrate (λ) ist der einzige Verteilungsparameter.

Inputanforderungen:

Rate von Lambda > 0

Lage kann jeder beliebiger positiver oder negativer Wert, einschließlich Null, sein

Extremwert- oder Gumbel-Verteilung

Die Extremwertverteilung (Typ 1) wird üblicherweise verwendet, um den größten Wert einer Antwort über einen Zeitraum zu beschreiben, zum Beispiel für Flutströme, Niederschläge und Erdbeben. Andere Anwendungen schließen Bruchkräfte von Materialien, Bauentwürfe und

Flugzeugbelastungen und -toleranzen ein. Die Extremwertverteilung ist auch als die Gumbel-Verteilung bekannt.

Die mathematischen Konstrukte für die Extremwertverteilung sind wie folgt:

$$f(x) = \frac{1}{\beta} z e^{-z} \text{ wobei } z = e^{\frac{x-\alpha}{\beta}} \text{ für } \beta > 0; \text{ und alle Werte von } x \text{ und } \alpha$$

$$\text{Mittelwert} = \alpha + 0.577215\beta$$

$$\text{Standardabweichung} = \sqrt{\frac{1}{6} \pi^2 \beta^2}$$

$$\text{Schiefe} = \frac{12\sqrt{6}(1.2020569)}{\pi^3} = 1.13955 \text{ (dies gilt für alle Werte des Modalwerts und der Skala)}$$

$$\text{Exzesswölbung} = 5.4 \text{ (dies gilt für alle Werte des Modalwerts und der Skala)}$$

Modalwert (α) und Skala (β) sind die Verteilungsparameter.

Berechnung der Parameter

Die Extremwertverteilung hat zwei Standardparameter: Modalwert und Skala. Der Modalwertparameter ist der wahrscheinlichste Wert für die Variable (der höchste Punkt auf der Wahrscheinlichkeitsverteilung). Nachdem Sie den Modalwertparameter ausgewählt haben, können Sie den Skalaparameter schätzen. Der Skalaparameter ist eine Zahl größer als 0. Je größer der Skalaparameter, umso größer die Varianz.

Inputanforderungen:

Modalwert Alpha kann alle Werte sein

Skala Beta > 0

F-Verteilung oder Fisher-Snedecor-Verteilung

Die F-Verteilung, auch als die Fisher-Snedecor-Verteilung bekannt, ist ebenfalls eine andere kontinuierliche Verteilung, die am häufigsten für Hypothesentesten verwendet wird. Im Besonderen wird sie verwendet, um den statistischen Unterschied zwischen zwei Varianzen in Analysen von Varianztests und Likelihood-Verhältnistests zu testen. Die F-Verteilung mit dem Freiheitsgradezähler n und dem Freiheitsgradenennern m ist mit der Chi-Quadrat Verteilung verwandt, insofern als:

$$\frac{\chi_n^2 / n}{\chi_m^2 / m} \sim F_{n,m}$$

$$\text{Mittelwert} = \frac{m}{m-2}$$

$$\text{Standardabweichung} = \frac{2m^2(m+n-2)}{n(m-2)^2(m-4)} \text{ für alle } m > 4$$

$$\text{Schiefe} = \frac{2(m+2n-2)}{m-6} \sqrt{\frac{2(m-4)}{n(m+n-2)}}$$

$$\text{Exzesswölbung} = \frac{12(-16+20m-8m^2+m^3+44n-32mn+5m^2n-22n^2+5mn^2)}{n(m-6)(m-8)(n+m-2)}$$

Der Freiheitsgradenzähler n und der Freiheitsgradenennner m sind die einzigen Verteilungsparameter.

Inputanforderungen:

Freiheitsgradenzähler und Freiheitsgradenennner beide > 0 Ganzzahlen

Gammaverteilung (Erlang-Verteilung)

Die Gammaverteilung ist in einem breiten Bereich von physikalischen Mengen anwendbar und ist mit anderen Verteilungen verwandt: Lognormal, Exponential, Pascal, Erlang, Poisson und Chi-Quadrat. Sie wird bei meteorologischen Prozessen verwendet, um Schadstoffkonzentrationen und Niederschlagsmengen zu repräsentieren. Die Gammaverteilung wird auch verwendet, um die Zeit zwischen den Vorkommen von Ereignissen zu messen, wenn der Ereignisprozess nicht komplett zufällig ist. Andere Anwendungen der Gammaverteilung schließen die Lagerbestandsführung, die Wirtschaftstheorie und die Versicherungsrisikothorie ein.

Bedingungen

Die Gammaverteilung wird meistens verwendet, als die Verteilung der Dauer bis das r -te Vorkommen eines Ereignisses in einem Poissonprozess. Wenn sie auf dieser Weise verwendet wird, sind die drei unterliegenden Bedingungen der Gammaverteilung:

- Die Anzahl der möglichen Vorkommen in einer Maßeinheit ist nicht auf eine feste Zahl begrenzt.
- Die Vorkommen sind unabhängig. Die Anzahl der Vorkommen in einer Maßeinheit beeinflusst nicht die Anzahl der Vorkommen in anderen Einheiten.
- Die Durchschnittsanzahl der Vorkommen muss von Einheit zu Einheit gleich bleiben.

Die mathematischen Konstrukte für die Gammaverteilung sind wie folgt:

$$f(x) = \frac{\left(\frac{x}{\beta}\right)^{\alpha-1} e^{-\frac{x}{\beta}}}{\Gamma(\alpha)\beta} \quad \text{mit allen Werten von } \alpha > 0 \text{ und } \beta > 0$$

$$\text{Mittelwert} = \alpha\beta$$

$$\text{Standardabweichung} = \sqrt{\alpha\beta^2}$$

$$\text{Schiefe} = \frac{2}{\sqrt{\alpha}}$$

$$\text{Exzesswölbung} = \frac{6}{\alpha}$$

Der Formparameter Alpha (α) und der Skalaparameter Beta (β) sind die Verteilungsparameter und Γ ist die Gammafunktion.

Wenn der Alphaparameter eine positive Ganzzahl ist, wird die Gammaverteilung als die Erlang-Verteilung bezeichnet. Sie wird verwendet, um die Wartezeiten in Warteschlangensystemen vorzuberechnen. Die Erlang-Verteilung ist die Summe von unabhängigen und identisch verteilten Zufallsvariablen, von denen jede eine gedächtnislose Exponentialverteilung hat. Wenn man n als die Anzahl dieser Zufallsvariablen einstellt, ist das mathematische Konstrukt der Erlang-Verteilung:

$$f(x) = \frac{x^{n-1} e^{-x}}{(n-1)!} \text{ für alle } x > 0 \text{ und alle positive Ganzzahlen von } n$$

Inputanforderungen:

Skala Beta > 0 und kann irgendein positiver Wert sein

Form Alpha ≥ 0.05 und alle positive Werte

Position kann irgendeinen Wert sein

Kosinus-Verteilung

Die Kosinusverteilung sieht wie eine logistische Verteilung aus, wobei der Medianwert zwischen Minimum und Maximum die höchste Spitze oder Modus hat, was die maximale Eintrittswahrscheinlichkeit mit sich trägt, während die extremen Schwänze, die nahe bei den minimalen und maximalen Werten liegen, niedrigere Wahrscheinlichkeiten haben. Minimum und Maximum sind die Verteilungsparameter.

Die mathematischen Konstrukte für die kosinus sind wie folgt:

$$f(x) = \begin{cases} \frac{1}{2b} \cos\left[\frac{x-a}{b}\right] & \text{for } \min \leq x \leq \max \\ 0 & \text{otherwise} \end{cases}$$

$$\text{where } a = \frac{\min + \max}{2} \text{ and } b = \frac{\max - \min}{\pi}$$

$$F(x) = \begin{cases} \frac{1}{2} \left[1 + \sin\left(\frac{x-a}{b}\right) \right] & \text{for } \min \leq x \leq \max \\ 1 & \text{for } x > \max \end{cases}$$

$$\text{Mittelwert} = \frac{\text{Min} + \text{Max}}{2}$$

$$\text{Standardabweichung} = \sqrt{\frac{(\text{Max} - \text{Min})^2 (\pi^2 - 8)}{4\pi^2}}$$

$$\text{Schiefe} = 0$$

$$\text{Exzesswölbung} = \frac{6(90 - \pi^4)}{5(\pi^2 - 6)^2}$$

Inputvoraussetzungen:

Minimum < Maximum

Laplace-Verteilung

Die Laplace-Verteilung wird mitunter auch die doppelte Exponentialverteilung genannt, denn man kann sie aus zwei nacheinander verbundenen Exponentialverteilungen zusammenstellen (mit einem zusätzlichen Lageparameter), was eine ungewöhnliche Spitze in der Mitte bildet. Die Wahrscheinlichkeitsdichtefunktion der Laplace-Verteilung erinnert an die Normalverteilung. Allerdings, während die Normalverteilung bezogen auf die quadratische Differenz vom Mittelwert ausgedrückt wird, wird die Laplace-Dichte bezogen auf die absolute Differenz vom Mittelwert ausgedrückt. Daraus folgt, dass die Schwänze der Laplace-Verteilung dicker als die der Normalverteilung sind. Wenn der Lageparameter auf Null eingestellt ist, wird die Zufallsvariable der Laplace-Verteilung exponentiell verteilt, mit einer Inversen des Skalaparameters. Alpha (auch als Lage bekannt) und Beta (auch als Skala bekannt) sind die Verteilungsparameter.

Die mathematischen Konstrukte für die Laplace sind wie folgt:

$$f(x) = \frac{1}{2\beta} \exp\left(-\frac{|x-\alpha|}{\beta}\right)$$
$$F(x) = \begin{cases} \frac{1}{2} \exp\left[\frac{x-\alpha}{\beta}\right] & \text{when } x < \alpha \\ 1 - \frac{1}{2} \exp\left[-\frac{x-\alpha}{\beta}\right] & \text{when } x \geq \alpha \end{cases}$$

Mittelwert = α

Standardabweichung = 1.4142β

Schiefe = 0

Exzesswölbung = 3

Inputvoraussetzungen:

Alpha (Lage) kann jeder beliebiger positiver oder negativer Wert, einschließlich Null, sein

Beta (Skala) > 0

Logistische Verteilung

Die logistische Verteilung ist gebräuchlich, um Wachstum zu beschreiben, das heißt, die Größe einer Bevölkerung ausgedrückt als Funktion einer Zeitvariablen. Sie kann auch verwendet werden, um chemische Reaktionen und den Wachstumskurs einer Bevölkerung oder eines Individuums zu beschreiben.

Die mathematischen Konstrukte für die logistische Verteilung sind wie folgt:

$$f(x) = \frac{e^{-\frac{\alpha-x}{\beta}}}{\beta \left[1 + e^{-\frac{\alpha-x}{\beta}} \right]^2} \quad \text{for alle Werte von } \alpha \text{ und } \mu$$

Mittelwert = α

$$\text{Standardabweichung} = \sqrt{\frac{1}{3} \pi^2 \beta^2}$$

Schiefe = 0 (dies gilt für alle Mittelwert- und Skala Inputs)

Exzesswölbung = 1.2 (dies gilt für alle Mittelwert- und Skala-Inputs)

Mittelwert (α) und Skala (β) sind die Verteilungsparameter.

Berechnung der Parameter

Die logistische Verteilung hat zwei Standardparameter: Mittelwert und Skala. Der Mittelwertparameter ist der Durchschnittswert, welcher für diese Verteilung der gleiche als der Modalwert ist, weil dies eine symmetrische Verteilung ist. Nachdem Sie den Modalwertparameter ausgewählt haben, können Sie den Skalaparameter schätzen. Der Skalaparameter ist eine Zahl größer als 0. Je größer der Skalaparameter, umso größer die Varianz. Inputanforderungen:

Skala Beta > 0 und kann alle positive Werte sein

Mittelwert Alpha kann irgendein Wert sein

Lognormalverteilung

Die Lognormalverteilung wird verbreitet in Situationen verwendet, wo Werte positiv asymmetrisch sind, zum Beispiel, in der Finanzanalyse für die Wertpapierbewertung oder im Immobiliensektor für die Immobilienbewertung und da wo Werte nicht unter Null fallen können. Aktienpreise sind normalerweise positiv asymmetrisch anstatt normal (symmetrisch) verteilt. Aktienpreise weisen diesen Trend auf, weil sie nicht unter die Grenze von Null fallen können, aber bis auf irgendeinem Preis ohne Limit steigen könnten. Ähnlicherweise stellen Immobilienpreise eine positive Schiefe dar, da Immobilienwerte nicht negativ werden können.

Bedingungen

Die drei unterliegenden Bedingungen der Lognormalverteilung sind:

- Die ungewisse Variable kann ohne Grenzen steigen, kann aber nicht unter Null fallen.
- Die ungewisse Variable ist positiv asymmetrisch, mit den meisten Werten in der Nähe der unteren Grenze.
- Der natürliche Logarithmus der ungewissen Variablen ergibt eine Normalverteilung.

Im Allgemeinen, wenn der Variabilitätskoeffizient größer als 30 Prozent ist, verwenden Sie eine Lognormalverteilung. Ansonsten verwenden Sie eine Normalverteilung.

Die mathematischen Konstrukte für die Lognormalverteilung sind wie folgt:

$$f(x) = \frac{1}{x\sqrt{2\pi} \ln(\sigma)} e^{-\frac{[\ln(x) - \ln(\mu)]^2}{2[\ln(\sigma)]^2}} \quad \text{für } x > 0; \mu > 0 \text{ und } \sigma > 0$$

$$\text{Mittelwert} = \exp\left(\mu + \frac{\sigma^2}{2}\right)$$

$$\text{Standardabweichung} = \sqrt{\exp(\sigma^2 + 2\mu)[\exp(\sigma^2) - 1]}$$

$$\text{Schiefe} = \left[\sqrt{\exp(\sigma^2) - 1}\right](2 + \exp(\sigma^2))$$

$$\text{Exzesswölbung} = \exp(4\sigma^2) + 2\exp(3\sigma^2) + 3\exp(2\sigma^2) - 6$$

Mittelwert (μ) und Standardabweichung (σ) sind die Verteilungsparameter.

Inputanforderungen:

Mittelwert und Standardabweichung beide > 0 und können irgendeinen positiven Wert haben

Lognormal Parametersätze

Die Lognormalverteilung verwendet standardmäßig den arithmetischen Mittelwert und die Standardabweichung. Für Anwendungen für die historische Daten verfügbar sind, ist es angemessener entweder den logarithmischen Mittelwert und die Standardabweichung, oder den geometrischen Mittelwert und die Standardabweichung zu verwenden.

Lognormale Versetzte Verteilung

Die Lognormalverteilung wird häufig verwendet in Situationen, wo Werte positiv asymmetrisch sind, zum Beispiel bei der Finanzanalyse für Kautionsbewertungen oder im Immobiliensektor für Immobilienbewertungen, und wo Werte nicht unter Null fallen können. Aktienpreise sind normalerweise positiv asymmetrisch anstatt normal (symmetrisch) verteilt. Aktienpreise weisen diesen Trend auf, weil sie nicht unter der unteren Grenze von Null fallen können, aber ohne Limit steigen könnten. Im Gegensatz dazu, die Lognormale versetzte Verteilung ist genau wie die Lognormalverteilung, aber so versetzt, dass der resultierende Wert auch negative Werte nehmen kann. Mittelwert (Durchschnittswert), Standardabweichung und Versetzung sind die Verteilungsparameter.

Inputvoraussetzungen:

Mittelwert > 0

Standardabweichung > 0

Versetzung kann jeder beliebiger positiver oder negativer Wert, einschließlich Null, sein

Normalverteilung

Die Normalverteilung ist die wichtigste Verteilung in der Wahrscheinlichkeitstheorie, weil sie

viele Naturphänomene, sowie den Intelligenzquotienten oder die Körpergröße von Menschen beschreibt. Entscheidungsträger können die Normalverteilung verwenden, um ungewisse Variablen, sowie die Inflationsrate oder den zukünftigen Preis von Benzin zu beschreiben.

Bedingungen

Die drei unterliegenden Bedingungen der Normalverteilung sind:

- Ein bestimmter Wert der ungewissen Variablen ist der wahrscheinlichste (der Mittelwert der Verteilung).
- Die ungewisse Variable könnte sowohl unter als auch über dem Mittelwert liegen (symmetrisch um den Mittelwert).
- Die ungewisse Variable liegt eher in der Nähe als weiter entfernt vom Mittelwert.

Die mathematischen Konstrukte für die Normalverteilung sind wie folgt:

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad \text{für alle Werte von } x \text{ und } \mu, \text{ während } \sigma > 0$$

Mittelwert = μ

Standardabweichung = σ

Schiefe = 0 (dies gilt für alle Mittelwert- und Standardabweichung-Inputs)

Exzesswölbung = 0 (dies gilt für alle Mittelwert- und Standardabweichung-Inputs)

Mittelwert (μ) und Standardabweichung (σ) sind die Verteilungsparameter.

Inputanforderungen:

Standardabweichung > 0 und kann alle positive Werte sein

Mittelwert kann ein beliebiger Wert sein

Parabolische-Verteilung

Die parabolische Austeilung ein Sonderfall von der Beta Austeilung ist, wenn Gestalt = maßstabgetreu Zeichnet = 2. Werte schließen zum Minimum und das Maximum hat niedrige Wahrscheinlichkeiten des Vorkommens, während Werte zwischen diesen zwei Extremen höhere Wahrscheinlichkeiten oder Vorkommen hat. Minimum und Maximum sind die distributional Parameter.

Die mathematischen Konstrukte für die Parabolische sind wie folgt:

$$f(x) = \frac{(x)^{(\alpha-1)}(1-x)^{(\beta-1)}}{\left[\frac{\Gamma(\alpha)\Gamma(\beta)}{\Gamma(\alpha+\beta)} \right]} \quad \text{for } \alpha > 0; \beta > 0; x > 0$$

Alpha = Beta = 2

Mittelwert = $\frac{Min + Max}{2}$

$$\text{Standardabweichung} = \sqrt{\frac{(\text{Max} - \text{Min})^2}{20}}$$

Schiefe = 0

Exzesswölbung = -0.8571

Inputvoraussetzungen:

Minimum < Maximum

Pareto-Verteilung

Die Pareto-Verteilung wird verbreitet für die Untersuchung von Verteilungen verwendet, die mit solchen empirischen Phänomenen assoziiert sind, wie die Größe von Stadtbevölkerungen, das Auftreten von Bodenschätzen, die Größe von Unternehmen, Personaleinkommen, Aktienpreisschwankungen und Fehlerclustering in Kommunikationskreislinien.

Die mathematischen Konstrukte für die Pareto sind wie folgt:

$$f(x) = \frac{\beta L^\beta}{x^{(1+\beta)}} \quad \text{for } x > L$$

$$\text{Mittelwert} = \frac{\beta L}{\beta - 1}$$

$$\text{Standardabweichung} = \sqrt{\frac{\beta L^2}{(\beta - 1)^2 (\beta - 2)}}$$

$$\text{Schiefe} = \sqrt{\frac{\beta - 2}{\beta} \left[\frac{2(\beta + 1)}{\beta - 3} \right]}$$

$$\text{Exzesswölbung} = \frac{6(\beta^3 + \beta^2 - 6\beta - 2)}{\beta(\beta - 3)(\beta - 4)}$$

Form (α) und Position (β) sind die Verteilungsparameter.

Berechnung der Parameter

Die Pareto-Verteilung hat zwei Standardparameter: Position und Form. Der Positionsparameter ist die untere Grenze für die Variable. Nachdem Sie den Positionsparameter ausgewählt haben, können Sie den Formparameter schätzen. Der Formparameter ist eine Zahl größer als 0, gewöhnlich größer als 1. Je größer der Formparameter, umso kleiner die Varianz und umso dicker der rechte Schwanz der Verteilung.

Inputanforderungen:

Position > 0 und kann alle positiven Werte sein

Form ≥ 0.05

Pearson-V-Verteilung

Die Pearson-V-Verteilung ist mit der inversen Gammaverteilung verwandt, wobei sie die Reziproke der nach der Gammaverteilung verteilten Variablen ist. Die Pearson-V-Verteilung wird auch verwendet, um Zeitverzögerungen zu modellieren, wenn es eine nahezu Gewissheit einer Minimalverzögerung gibt und die Maximalverzögerung unbegrenzt ist, z.B., die Verzögerung bei der Ankunft von Notdiensten und die erforderliche Zeit, um eine Maschine zu reparieren. Alpha (auch als Form bekannt) und Beta (auch als Skala bekannt) sind die Verteilungsparameter.

Die mathematischen Konstrukte für die Pearson V sind wie folgt:

$$f(x) = \frac{x^{-(\alpha+1)} e^{-\beta/x}}{\beta^{-\alpha} \Gamma(\alpha)}$$

$$F(x) = \frac{\Gamma(\alpha, \beta/x)}{\Gamma(\alpha)}$$

$$\text{Mittelwert} = \frac{\beta}{\alpha - 1}$$

$$\text{Standardabweichung} = \sqrt{\frac{\beta^2}{(\alpha - 1)^2 (\alpha - 2)}}$$

$$\text{Schiefe} = \frac{4\sqrt{\alpha - 2}}{\alpha - 3}$$

$$\text{Exzesswölbung} = \frac{30\alpha - 66}{(\alpha - 3)(\alpha - 4)} - 3$$

Inputvoraussetzungen:

Alpha (Form) > 0

Beta (Skala) > 0

Pearson-VI-Verteilung

Die Pearson-VI-Verteilung ist mit der Gammaverteilung verwandt, wobei sie die rationale Funktion von zwei nach der Gammaverteilung verteilten Variablen ist. Alpha 1 (auch als Form 1 bekannt), Alpha 2 (auch als Form 2 bekannt) und Beta (auch als Skala bekannt) sind die Verteilungsparameter.

Die mathematischen Konstrukte für die Pearson VI sind wie folgt:

$$f(x) = \frac{(x/\beta)^{\alpha_1-1}}{\beta B(\alpha_1, \alpha_2) [1 + (x/\beta)]^{\alpha_1+\alpha_2}}$$

$$F(x) = F_B\left(\frac{x}{x+\beta}\right)$$

$$\text{Mittelwert} = \frac{\beta\alpha_1}{\alpha_2 - 1}$$

$$\text{Standardabweichung} = \sqrt{\frac{\beta^2\alpha_1(\alpha_1 + \alpha_2 - 1)}{(\alpha_2 - 1)^2(\alpha_2 - 2)}}$$

$$\text{Schiefe} = 2\sqrt{\frac{\alpha_2 - 2}{\alpha_1(\alpha_1 + \alpha_2 - 1)}} \left[\frac{2\alpha_1 + \alpha_2 - 1}{\alpha_2 - 3} \right]$$

$$\text{Exzesswölbung} = \frac{3(\alpha_2 - 2)}{(\alpha_2 - 3)(\alpha_2 - 4)} \left[\frac{2(\alpha_2 - 1)^2}{\alpha_1(\alpha_1 + \alpha_2 - 1)} + (\alpha_2 + 5) \right] - 3$$

Inputvoraussetzungen:

Alpha 1 (Form 1) > 0

Alpha 2 (Form 2) > 0

Beta (Skala) > 0

PERT-Verteilung

Die PERT-Verteilung wird häufig in Projekt- und Programmmanagement verwendet, um die Szenarien des schlimmsten, des nominalen und des besten Falls der Projektfertigstellungszeit zu bestimmen. Sie ist mit den Beta- und Dreiecksverteilungen verwandt. Man kann die PERT-Verteilung verwenden, um Risiken in Projekt- und Kostenmodellen zu identifizieren, basierend auf der Wahrscheinlichkeit der Erfüllung von Vorgaben und Zielen einer beliebigen Anzahl von Projektkomponenten und unter Verwendung von minimalen, am wahrscheinlichsten und maximalen Werten. Sie ist aber auch konzipiert, um eine Verteilung zu generieren, die Wahrscheinlichkeitsverteilungen am meisten ähnelt. Die PERT-Verteilung kann eine nahe Anpassung an Normal- oder Lognormalverteilungen bieten. Wie die Dreiecksverteilung betont die PERT-Verteilung den "am wahrscheinlichsten" Wert anstelle der Minimal- oder Maximalschätzungen. Allerdings, anders als die Dreiecksverteilung, bildet die PERT-Verteilung eine glatte Kurve, die zunehmend mehr Betonung auf Werte um (in der Nähe) den am wahrscheinlichsten Wert setzt, statt auf Werte um die Grenzen. Konkret heißt das, dass wir die Schätzung des am wahrscheinlichsten Werts "trauen", und wir glauben, auch wenn sie nicht exakt genau ist (wie Schätzungen es selten sind), dass wir erwarten können, dass der resultierende Wert nahe dieser Schätzung sein wird. Wenn man annimmt, dass viele echte Welt Phänomenen normal verteilt sind, liegt die Attraktivität der PERT-Verteilung darin, dass sie eine ähnliche Kurve in Form wie die Normalkurve produziert, ohne dass man die genauen Parameter der verwandten Normalkurve kennen muss. Minimum, am wahrscheinlichsten und Maximum sind die Verteilungsparameter.

Die mathematischen Konstrukte für die PERT sind wie folgt:

$$f(x) = \frac{(x - \min)^{A1-1} (\max - x)^{A2-1}}{B(A1, A2)(\max - \min)^{A1+A2-1}}$$

$$\text{where } A1 = 6 \left[\frac{\min + 4(\text{likely}) + \max}{\max - \min} - \min \right] \text{ and } A2 = 6 \left[\max - \frac{\min + 4(\text{likely}) + \max}{\max - \min} \right]$$

and B is the Beta function

$$\text{Mittelwert} = \frac{\text{Min} + 4\text{Mode} + \text{Max}}{6}$$

$$\text{Standardabweichung} = \sqrt{\frac{(\mu - \text{Min})(\text{Max} - \mu)}{7}}$$

$$\text{Schiefe} = \sqrt{\frac{7}{(\mu - \text{Min})(\text{Max} - \mu)}} \left(\frac{\text{Min} + \text{Max} - 2\mu}{4} \right)$$

Inputvoraussetzungen:

Minimum $\leq$ am wahrscheinlichsten $\leq$ Maximum und kann positiv, negativ oder Null sein

Potenzverteilung

Die Potenzverteilung ist mit der Exponentialverteilung insofern verwandt, dass die Wahrscheinlichkeit von kleinen Ausgängen groß ist, aber mit der Steigerung des Ausgangswerts vermindert sich diese exponentiell. Alpha (auch als Form bekannt) ist der einzige Verteilungsparameter.

Die mathematischen Konstrukte für die potenz sind wie folgt:

$$f(x) = \alpha x^{\alpha-1}$$

$$F(x) = x^{\alpha}$$

$$\text{Mittelwert} = \frac{\alpha}{1 + \alpha}$$

$$\text{Standardabweichung} = \sqrt{\frac{\alpha}{(1 + \alpha)^2 (2 + \alpha)}}$$

$$\text{Schiefe} = \sqrt{\frac{\alpha + 2}{\alpha}} \left(\frac{2(\alpha - 1)}{\alpha + 3} \right)$$

Inputvoraussetzungen:

Alpha (Form) > 0

Multiplikative Potenzverteilung versetzte

Die Potenzverteilung ist mit der Exponentialverteilung insofern verwandt, dass die Wahrscheinlichkeit von kleinen Ausgängen groß ist, aber mit der Steigerung des Ausgangswerts vermindert sich diese exponentiell. Alpha (auch als Form bekannt) ist der einzige Verteilungsparameter.

Inputvoraussetzungen:

Alpha (Form) > 0

Lage kann jeder beliebiger positiver oder negativer Wert, einschließlich Null, sein

Faktor > 0

Studentsche t Verteilung

Die Studentsche t-Verteilung ist die beim Hypothesentesten am meisten verwendete Verteilung. Diese Verteilung wird verwendet, um den Mittelwert einer normalverteilten Bevölkerung zu schätzen, wenn die Stichprobengröße klein ist. Sie wird verwendet, um die statistische Signifikanz des Unterschieds zwischen zwei Stichprobenmittelwerten oder Konfidenzintervallen für kleine Stichprobengrößen zu testen.

Die mathematischen Konstrukte für die t-Verteilung sind wie folgt:

$$f(t) = \frac{\Gamma[(r+1)/2]}{\sqrt{r\pi} \Gamma[r/2]} (1 + t^2/r)^{-(r+1)/2}$$

Mittelwert = 0 (dies gilt für alle Freiheitsgrade r , außer wenn die Verteilung zu einer anderen Nichtnull Zentralposition verschoben wird)

$$\text{Standardabweichung} = \sqrt{\frac{r}{r-2}}$$

Schiefe = 0 (dies gilt für alle Freiheitsgrade r)

$$\text{Exzesswölbung} = \frac{6}{r-4} \text{ für alle } r > 4$$

wobei $t = \frac{x - \bar{x}}{s}$ und Γ die Gammafunktion ist.

Freiheitsgrade r sind die einzigen Verteilungsparameter.

Die t-Verteilung ist mit der F-Verteilung wie folgt verwandt: die Quadratzahl eines Wertes von t mit r Freiheitsgraden ist als F mit 1 und r Freiheitsgraden verteilt. Die Gesamtform der Wahrscheinlichkeitsdichtefunktion der t-Verteilung ähnelt auch der Glockenform einer normalverteilten Variablen mit Mittelwert 0 und Varianz 1, außer dass sie ein bisschen niedriger und breiter oder leptokurtisch ist (dicke Schwänze an den Enden und spitzes Zentrum). Mit dem Wachsen der Anzahl der Freiheitsgrade (sagen wir über 30), nähert sich die t-Verteilung der Normalverteilung mit Mittelwert 0 und Varianz 1 an.

Inputanforderungen:

Freiheitsgrade ≥ 1 und muss eine Ganzzahl sein

Uniformverteilung

Bei der Uniformverteilung ereignen sich alle Werte, die zwischen Minimum und Maximum fallen, mit der gleichen Wahrscheinlichkeit.

Bedingungen

Die drei unterliegenden Bedingungen der Uniformverteilung sind:

- Der Minimalwert ist festgelegt.
- Der Maximalwert ist festgelegt..
- Alle Werte zwischen Minimum und Maximum ereignen sich mit der gleichen Wahrscheinlichkeit.

Die mathematischen Konstrukte für die Uniformverteilung sind wie folgt:

$$f(x) = \frac{1}{Max - Min} \quad \text{für alle Werte, sodass } Min < Max$$

$$\text{Mittelwert} = \frac{Min + Max}{2}$$

$$\text{Standardabweichung} = \sqrt{\frac{(Max - Min)^2}{12}}$$

Schiefe = 0 (dies gilt für alle Inputs von *Min* und *Max*)

Exzesswölbung = -1.2 (dies gilt für alle Inputs von *Min* und *Max*)

Maximalwert (*Max*) und Minimalwert (*Min*) sind die Verteilungsparameter.

Inputanforderungen:

$Min < Max$ und kann alle Werte sein

Weibull-Verteilung (Rayleigh-Verteilung)

Die Weibull-Verteilung beschreibt Daten, die von Lebensdauer- und Ermüdungstests stammen. Sie wird üblicherweise verwendet, um die Ausfallszeit in Zuverlässigkeitsstudien und die Bruchstärken von Materialien in Zuverlässigkeits- und Qualitätssicherungstests zu beschreiben. Weibull-Verteilungen werden auch verwendet, um verschiedene physikalische Mengen, wie die Windgeschwindigkeit, zu repräsentieren.

Die Weibull-Verteilung ist eine Familie von Verteilungen, welche die Eigenschaften von mehreren anderen Verteilungen annehmen kann. Zum Beispiel, abhängig von dem von Ihnen definierten Formparameter, kann man die Weibull-Verteilung verwenden, um, unter anderem, die Exponential- und Rayleigh-Verteilungen zu modellieren. Die Weibull-Verteilung ist sehr flexibel. Wenn der Weibull-Formparameter 1.0 gleicht, ist die Weibull-Verteilung identisch mit der Exponentialverteilung. Der Weibull-Positionsparameter erlaubt es Ihnen eine Exponentialverteilung aufzustellen, die bei einer anderen Position als 0.0 beginnen soll. Wenn der Formparameter kleiner als 1.0 ist, wird die Weibull-Verteilung eine stark abfallende Kurve. Ein Hersteller könnte diesen Effekt bei der Beschreibung von Bauteilausfällen während einer Einbrennphase finden.

Die mathematischen Konstrukte für die Weibull-Verteilung sind wie folgt:

$$f(x) = \frac{\alpha}{\beta} \left[\frac{x}{\beta} \right]^{\alpha-1} e^{-\left(\frac{x}{\beta}\right)^{\alpha}}$$

$$\text{Mittelwert} = \beta \Gamma(1 + \alpha^{-1})$$

$$\text{Standardabweichung} = \beta^2 [\Gamma(1 + 2\alpha^{-1}) - \Gamma^2(1 + \alpha^{-1})]$$

$$\text{Schiefe} = \frac{2\Gamma^3(1 + \beta^{-1}) - 3\Gamma(1 + \beta^{-1})\Gamma(1 + 2\beta^{-1}) + \Gamma(1 + 3\beta^{-1})}{[\Gamma(1 + 2\beta^{-1}) - \Gamma^2(1 + \beta^{-1})]^{3/2}}$$

$$\text{Exzesswölbung} =$$

$$\frac{-6\Gamma^4(1 + \beta^{-1}) + 12\Gamma^2(1 + \beta^{-1})\Gamma(1 + 2\beta^{-1}) - 3\Gamma^2(1 + 2\beta^{-1}) - 4\Gamma(1 + \beta^{-1})\Gamma(1 + 3\beta^{-1}) + \Gamma(1 + 4\beta^{-1})}{[\Gamma(1 + 2\beta^{-1}) - \Gamma^2(1 + \beta^{-1})]^2}$$

Form (α) und zentrale Positionsskala (β) sind die Verteilungsparameter und Γ ist die Gammafunktion.

Inputanforderungen:

Form Alpha ≥ 0.05

Skala Beta > 0 und kann alle positiven Werte sein

Multiplikative Weibull- und Rayleighverteilungen versetzte

Die Weibull-Verteilung beschreibt Daten, die von Lebensdauer- und Ermüdungstests stammen. Sie wird üblicherweise verwendet, um die Ausfallszeit in Zuverlässigkeitsstudien sowie die Bruchstärken von Materialien in Zuverlässigkeits- und Qualitätssicherungsprüfungen zu beschreiben. Weibull-Verteilungen werden auch verwendet, um verschiedene physikalische Mengen, sowie die Windgeschwindigkeit, zu repräsentieren. Die Weibull-Verteilung ist eine Familie von Verteilungen, welche die Eigenschaften von mehreren anderen Verteilungen annehmen kann. Zum Beispiel, abhängig von dem von Ihnen definierten Formparameter, kann man die Weibull-Verteilung verwenden, um, unter anderem, die Exponential- und Rayleighverteilungen zu modellieren. Die Weibull-Verteilung ist sehr flexibel. Wenn der Weibull-Formparameter 1.0 gleicht, ist die Weibull-Verteilung identisch mit der

Exponentialverteilung. Die Weibull zentrale Positionsskala oder der Betaparameter erlaubt es Ihnen eine Exponentialverteilung auszustellen, die bei einer anderen Position als 0.0 beginnen soll. Wenn der Formparameter kleiner als 1.0 ist, wird die Weibull-Verteilung eine stark abfallende Kurve. Ein Hersteller könnte diesen Effekt bei der Beschreibung von Bauteilausfällen während einer Einbrennphase nützlich finden. Form (α) und Skala (β) sind die Verteilungsparameter.

Inputvoraussetzungen:

Form Alpha ≥ 0.05

Zentrale Positionsskala oder Beta > 0 und kann jeden beliebigen Wert haben

Lage kann jeder beliebiger positiver oder negativer Wert, einschließlich Null, sein

Faktor > 0

3. VORAUSBERECHNUNG

Die Ausführung einer Vorausberechnung ist in der Tat die Prognostizierung der Zukunft, ob auf Basis von historischen Daten oder Mutmaßungen über die Zukunft, wenn keine Historie vorhanden ist. Wenn historischen Daten existieren, ist eine quantitative oder statistische Methode die beste; Wenn aber keine historischen Daten existieren, dann ist normalerweise eine qualitative oder verurteilende Methode eventuell die einzige Möglichkeit. Bild 3.1 listet die am gebräuchlichsten Methodologien der Vorausberechnung.

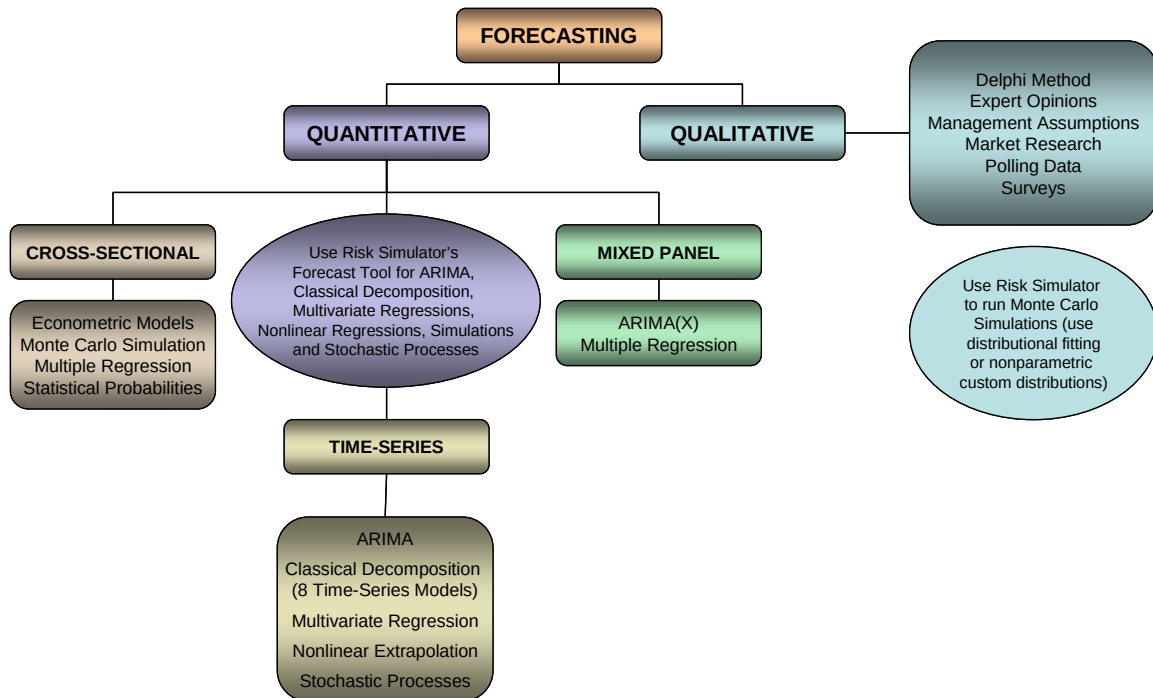


Bild 3.1—Vorausberechnungsmethoden

Verschiedenen Typen von Vorausberechnungsverfahren

Allgemein kann man die Vorausberechnung in quantitative und qualitative teilen. Die qualitative Vorausberechnung wird verwendet, wenn wenige oder keine historischen, gegenwärtigen oder vergleichbaren Daten vorhanden sind. Es existieren verschiedene qualitative Methoden: Delphi oder Expertenmeinung (eine Konsensbildende Vorausberechnung von Sektorenexperten, Marketingexperten oder internen Mitarbeitern); Managementhypothesen (vom Führungsstab festgelegte Wachstumszielraten); sowie auch Marktforschung, externe Daten oder Meinungsumfragen und -erhebungen (Daten, die von Fremdquellen, Industrie und Branchenindexen, oder von aktiver Marktforschung stammen). Diese Schätzungen können entweder Einzelpunkt-Schätzungen (ein durchschnittlicher Konsens) oder ein Satz von Vorausberechnungswerten (eine Verteilung von Vorausberechnungen) sein. Man kann diese

letzteren in *Risiko Simulator* als eine angepasste Verteilung eingeben und die resultierenden Vorausberechnungen simulieren. Das heißt, eine nicht-parametrische Simulation unter Verwendung der geschätzten Datenpunkte selber als die Verteilung.

Bei dem quantitativen Vorausberechnungstyp kann man die verfügbaren Daten oder die Daten, die vorauszuberechnen sind, folgendermaßen teilen: Zeitreihendaten (Werte, die ein Zeitelement haben, wie Einnahmen in verschiedenen Jahren, Inflationsraten, Zinssätze, Marktanteil, Ausfallraten und so weiter); Querschnittsdaten (Werte, die Zeitunabhängig sind, wie der Notendurchschnitt von Zehntklässlern bundesweit in einem bestimmten Jahr, gegeben die Punkte des SAT (Zulassungstest für amerikanischen Hochschulen), den IQ und die Anzahl der pro Woche konsumierten Alkoholika jedes Studenten); oder Mischpaneldaten (Mischung zwischen Zeitreihen- und Paneldaten, z.B., die Vorausberechnung von Verkäufen über den nächsten 10 Jahre, gegeben die vorgesehenen Marketingausgaben und Marktanteilprojektionen; dies bedeutet, dass die Verkaufsdaten Zeitreihen sind, aber dass exogene Variablen wie Marketingausgaben und Marktanteil vorhanden sind, um Hilfe bei der Modellierung von Vorausberechnungsvorhersagen zu leisten).

Die Software *Risiko Simulator* stellt dem Benutzer einige Vorausberechnungsmethodologien zu Verfügung:

1. ARIMA (Autoregressiver integrierter gleitender Mittelwert)
2. Auto-ARIMA
3. Auto-Ökonometrie
4. Grund-Ökonometrie
5. Angepasste Verteilungen
6. Kombinatorischer Fuzzylogik
7. GARCH (Verallgemeinerte autoregressive bedingte Heteroskedastizität)
8. J-Kurven
9. Markov-Ketten
10. Maximale Wahrscheinlichkeit (Logit, Probit, Tobit)
11. Neuronales Netzwerk
12. Multivariate Regression
13. Nichtlineare Extrapolation
14. S-Kurven
15. Kubischer Spline-Kurven
16. Stochastische Prozess Vorausberechnung
17. Analyse und Zerlegung von Zeitreihen
18. Trendlinien

Die analytischen Details von jeder Vorausberechnungsmethode fallen außerhalb der Materie dieses Benutzerhandbuchs. Für mehr Details, bitte lesen Sie *Modeling Risk*, 2nd Edition: *Applying Monte Carlo Simulation, Real Options Analysis, Stochastic Forecasting and Portfolio Optimization*, von Dr. Johnathan Mun (Wiley Finance 2010), welcher auch der Erschaffer der Software *Risiko Simulator* ist. Trotzdem, das Folgende erläutert einige der gängigsten Methoden.

Alle anderen Vorausberechnungsmethoden sind ziemlich einfach innerhalb Risiko Simulator anzuwenden. Das Folgende stellt eine schnelle Übersicht von jeder Methodologie und einige schnelle Beispiel zur Verfügung, um mit der Verwendung der Software zu beginnen. Mehr detaillierte Beschreibungen und Beispielsmodelle von jeder dieser Methoden sind in diesem ganzen und den nächsten Kapiteln zu finden.

② ARIMA

Der autoregressive integrierte gleitende Mittelwert (ARIMA, auch als Box-Jenkins ARIMA bekannt) ist ein fortgeschrittenes ökonometrisches Modellierungsverfahren. ARIMA examiniert historische Zeitreihendaten und führt Zurückanpassungs-Optierungsroutinen aus, um folgendes zu berücksichtigen: Die historische Autokorrelation (das Verhältnis von einem Wert zu einem anderen im Laufe der Zeit) und die Stabilität der Daten, um die Nichtstationaritätseigenschaften der Daten zu korrigieren. Außerdem lernt dieses prädiktive Modell im Laufe der Zeit, indem es seine Vorausberechnungsfehler korrigiert. Normalerweise ist eine fortgeschrittene Kenntnis der Ökonometrie erforderlich, um gute prädiktive Modelle unter Verwendung dieser Methode aufzubauen.

② Auto-ARIMA

Das Auto-ARIMA Modul automatisiert einiges in der herkömmlichen ARIMA-Modellierung, indem es automatisch mehrfache Permutationen von Modellspezifikationen testet und das bestpassende Modell zurückgibt. Die Ausführung von Auto-ARIMA ist ähnlich wie bei den normalen ARIMA-Vorausberechnungen. Der Unterschied ist, dass die P, D, Q Inputs nicht länger erforderlich sind und dass verschiedene Kombinationen dieser Inputs automatisch ausgeführt und verglichen werden.

② Grund-Ökonometrie

Die Ökonometrie bezieht sich auf eine Branche der Geschäftsanalytik: Modellierungs- und Vorausberechnungsverfahren zur Modellierung des Verhaltens und zur Vorausberechnung von bestimmten Variablen im Geschäftsleben, in der Wirtschaft, in der Finanz, in der Physik, in der Herstellung, im Betrieb und von allen anderen Variablen. Die Ausführung der Modelle der Grund-Ökonometrie ist ähnlich wie bei der normalen Regressionsanalyse, außer dass man die abhängigen und unabhängigen Variablen vor der Ausführung einer Regression modifizieren kann.

② Grund-Auto-Ökonometrie

Ähnlich wie die Grund-Ökonometrie, aber es werden Tausende von linearen, nichtlinearen, interagierenden, verzögerten und gemischten Variablen auf Ihre Daten automatisch ausgeführt, um das bestpassende ökonometrische Modell festzustellen, welches das Verhalten der abhängigen Variablen beschreibt. Dies ist nützlich für das Modellieren der Effekte von Variablen und für die Vorausberechnung von zukünftigen Ausgängen, ohne dass der Analyst unbedingt ein Experte Ökonometriker sein muss.

② Angepasste Verteilungen

Unter Verwendung des Risiko Simulators kann man Expertenmeinungen sammeln und eine angepasste Verteilung generieren. Dieses Vorausberechnungsverfahren erweist sich als nützlich, wenn

bei der Anwendung auf einer Verteilungsanpassungsroutine der Datensatz klein oder die Güte-der-Anpassung schlecht ist.

② GARCH

Das Modell der verallgemeinerten autoregressiven bedingten Heteroskedastizität (GARCH) wird verwendet, um von einem marktgängigen Wertpapier (z.B., Aktienpreise, Rohstoffpreise, Erdölpreise und so weiter) die historischen Volatilitätsniveaus zu modellieren und die zukünftigen Volatilitätsniveaus vorzuberechnen. Der Datensatz muss eine Zeitreihe von Rohpreisniveaus sein. Erst konvertiert GARCH die Preise in relative Erträge und dann wird eine interne Optimierung ausgeführt, um die historischen Daten in eine zum Mittelwert zurückkehrende Volatilitätsterminstruktur anzupassen, unter der Annahme, dass die Natur der Volatilität heteroskedastisch ist (sie ändert sich im Laufe der Zeit gemäß einiger ökonomischer Eigenschaften). Mehrere Variationen dieser Methodologie sind in Risiko Simulator verfügbar, einschließlich EGARCH, EGARCH-T, GARCH-M, GJR-GARCH, GJR-GARCH-T, IGARCH und T-GARCH.

② J-Kurve

Die J-Kurve oder exponentielle Wachstumskurve ist eine Kurve wo das Wachstum der nächsten Periode vom Niveau der aktuellen Periode abhängt und der Anstieg exponentiell ist. Das heißt, dass im Laufe der Zeit, die Werte von einer Periode zur anderen signifikant wachsen werden. Dieses Modell wird typisch in der Vorausberechnung des biologischen Wachstums und der chemischen Reaktionen im Laufe der Zeit verwendet.

② Markov-Ketten

Eine Markov-Kette existiert, wenn die Wahrscheinlichkeit eines zukünftigen Zustands von einem vorhergehenden Zustand abhängt und wenn zusammengefügt, sie eine Kette bilden, die zurück zu einem Langzeit-Dauerzustandsniveau kehren. Diese Methode wird typisch verwendet, um den Marktanteil von zwei Konkurrenten vorzuberechnen. Die erforderlichen Inputs sind die Anfangswahrscheinlichkeit, dass ein Kunde des ersten Geschäfts (der erste Zustand) zu diesem selben Geschäft in der nächsten Periode zurückkehren wird, gegen die Wahrscheinlichkeit, dass er in dem nächsten Zustand zum Geschäft eines Konkurrenten wechseln wird.

② Maximale Wahrscheinlichkeit auf Logit, Tobit und Probit

Die Schätzung der maximalen Wahrscheinlichkeit (MLE) wird verwendet, um die Wahrscheinlichkeit des Vorkommens eines Ereignisses vorzuberechnen, gegeben bestimmte unabhängige Variablen. MLE wird verwendet, zum Beispiel, um vorherzusagen, ob eine Kreditlinie oder Zahlungsverpflichtung in Verzug geraten wird, gegeben die Eigenschaften des Schuldners (30 Jahre alt, Alleinstehender, Gehalt von \$100.000 pro Jahr und eine Gesamtkreditkartenverschuldung von \$10.000). Ebenso kann man die Wahrscheinlichkeit vorzuberechnen, ob ein Patient Lungenkrebs entwickeln wird, wenn die Person ein Mann zwischen 50 und 60 Jahre alt ist, 5 Packungen Zigaretten pro Monat raucht und so weiter. Unter diesen Umständen ist die abhängige Variable begrenzt (das heißt, sie ist nur binär, 1 und 0 für Verzug/Sterben und kein Verzug/Leben, oder begrenzt auf Ganzzahlwerte wie 1, 2, 3 und so weiter) und der gewünschte Ausgang des Modells ist es, die

Wahrscheinlichkeit des Vorkommens eines Ereignisses vorherzusagen. Die herkömmliche Regressionsanalyse funktioniert nicht in diesen Situationen (die vorausberechnete Wahrscheinlichkeit ist normalerweise weniger als Null oder größer als 1, viele der erforderlichen Hypothesen, sowie auch die Unabhängigkeit und Normalität des Fehlers, werden verletzt und die Fehler werden ziemlich groß sein).

② Multivariate Regression

Die multivariate Regression wird verwendet, um die Verhältnisstruktur und Eigenschaften einer bestimmten abhängigen Variablen zu modellieren, da diese von anderen unabhängigen Variablen abhängt. Unter Verwendung des modellierten Verhältnisses kann man die zukünftigen Werte der abhängigen Variablen vorausberechnen. Man kann auch die Genauigkeit und Güte-der-Anpassung für dieses Modell bestimmen. Man kann lineare und nichtlineare Modelle in die mehrfache Regressionsanalyse anpassen.

② Nichtlineare Extrapolation

Es wird angenommen, dass die unterliegende Struktur der zu vorausberechnenden Daten nichtlinear im Laufe der Zeit ist. Zum Beispiel, ein Datensatz wie 1, 4, 9, 16, 25 gilt als nichtlinear (diese Datenpunkte stammen von einer Quadratfunktion).

② S-Kurven

Die S-Kurve oder logistische Wachstumskurve beginnt wie eine J-Kurve mit exponentiellen Wachstumsraten. Im Laufe der Zeit wird die Umgebung gesättigt (z.B., Marksättigung, Konkurrenz, Überfüllung), das Wachstum lässt nach und der Vorausberechnungswert endet schließlich am Sättigungs- oder Maximalniveau. Dieses Modell wird typischerweise verwendet bei der Vorausberechnung des Marktanteils oder des Verkaufswachstums eines neuen Produktes von der Markteinführung bis zur Reife und den Rückgang, bei der Bevölkerungsdynamik und bei anderen natürlich vorkommenden Phänomenen.

② Kubischer Spline-Kurven

Gelegentlich gibt es fehlende Werte in einem Zeitreihendatensatz. Man hat, zum Beispiel, die Zinssätze für die Jahre 1 bis 3, gefolgt von den Jahren 5 bis 8 und dann für Jahr 10. Man kann Spline-Kurven verwenden, um die Zinssatzwerte der fehlenden Jahre, basierend auf den existierenden Daten, zu interpolieren. Spline-Kurven können auch verwendet werden, um Werte von zukünftigen Perioden über die Zeitperiode der existierenden Daten hinaus vorauszuberechnen oder zu extrapolieren. Die Daten können linear oder nichtlinear sein.

② Stochastischer Prozess Vorausberechnung

Gelegentlich sind Variable stochastisch und man kann sie nicht leicht unter Verwendung von herkömmlichen Methoden vorausberechnen. Diese Variablen werden als stochastisch bezeichnet. Trotzdem folgen die meisten finanziellen, wirtschaftlichen und natürlich vorkommenden Phänomene (z.B., die Bewegung von Molekülen durch die Luft) einem bekannten mathematischen Gesetz oder Verhältnis. Obwohl die resultierenden Werte ungewiss sind, ist die unterliegende mathematische Struktur bekannt und kann unter Verwendung einer Monte-Carlo-Risikosimulation simuliert werden.

Die in Risiko Simulator gestützten Prozesse schließen die folgenden ein: Brownsche Bewegung Irrfahrt, Rückkehr zum Mittelwert, Sprung-Diffusion und gemischte Prozesse, nützlich für die Vorausberechnung von nichtstationären Zeitreihenvariablen.

☺ Analyse und Zerlegung von Zeitreihen

In Zeitreihendaten mit „gutem Verhalten“ (typische Beispiele schließen Verkaufseinnahmen und Kostenstrukturen von großen Unternehmen ein), haben die Werte meistens bis zu drei Elemente; Basiswert, Trend und Saisonalität. Die Zeitreihenanalyse verwendet diese historischen Daten, zerlegt sie in diese drei Elemente und setzt sie dann in Zukunftsvorausberechnungen wieder zusammen. In anderen Worten führt diese Vorausberechnungsmethode, sowie einige der anderen hier beschriebenen, erst eine Rückanpassung (Rückberechnung) der historischen Daten aus, bevor sie dann Schätzungen von Zukunftswerten (Vorausberechnungen) liefert.

Das Vorausberechnungs-Tool in Risiko Simulator ausführen

Generell, um Vorausberechnungen zu kreieren, sind einige schnelle Schritte erforderlich:

- Excel starten und auf Ihre existierenden historischen Daten zugreifen oder sie öffnen
- Die Daten auswählen, auf *Simulation* klicken und *Vorausberechnung* wählen
- Die zutreffenden Sektionen (ARIMA, Multivariate Regression, Nichtlineare Extrapolation, Stochastische Vorausberechnung, Zeitreihenanalyse) auswählen und die relevanten Inputs eingeben

Bild 3.2 zeigt das Tool *Vorausberechnung* und die verschiedenen Methodologien.



Bild 3.2—Vorausberechnungsmethoden von Risiko Simulator

Das Folgende stellt eine schnelle Übersicht jeder Methodologie und einige schnelle Beispiele, um mit der Verwendung der Software zu beginnen zur Verfügung. Man kann die Beispielsdatei entweder im Startmenü unter [Start](#) | [Real Options Valuation](#) | [Risiko Simulator](#) | [Beispiele](#) finden oder direkt durch [Risiko Simulator](#) | [Beispielsmodelle](#) aufrufen.

Zeitreihenanalyse

Theorie:

Bild 3.3 listet die 8 gängigsten Zeitreihenmodelle, getrennt nach Saisonalität und Trend. Zum Beispiel, wenn die Datenvariable keinen Trend oder keine Saisonalität hat, dann würde ein „*einzelner gleitender Mittelwert*“ Modell oder ein „*einzelne exponentielle Glättung*“ Modell genügen. Wenn jedoch die Saisonalität existiert, aber kein erkennbarer Trend vorhanden ist, dann wäre entweder ein Modell der additiven Saisonalität oder der multiplikativen Saisonalität besser, und so weiter.

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

Bild 3.3—Die 8 gängigsten Zeitreihenmethoden

Prozedur:

- Excel starten und bei Bedarf Ihre historischen Daten öffnen (das nachstehende Beispiel verwendet die Datei **Zeitreihen-Vorausberechnung** im Ordner Beispiele)
- Wählen Sie die historischen Daten aus (die Daten sollten in einer einzelnen Spalte aufgelistet werden)
- Wählen Sie **Risiko Simulator | Vorausberechnung | Zeitreihenanalyse**
- Wählen Sie das anzuwendenden Modell aus, geben Sie die relevanten Hypothesen ein und klicken Sie auf OK

Historische Verkaufseinnahmen

Jahr	Vierteljahr	Periode	Umsatz
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

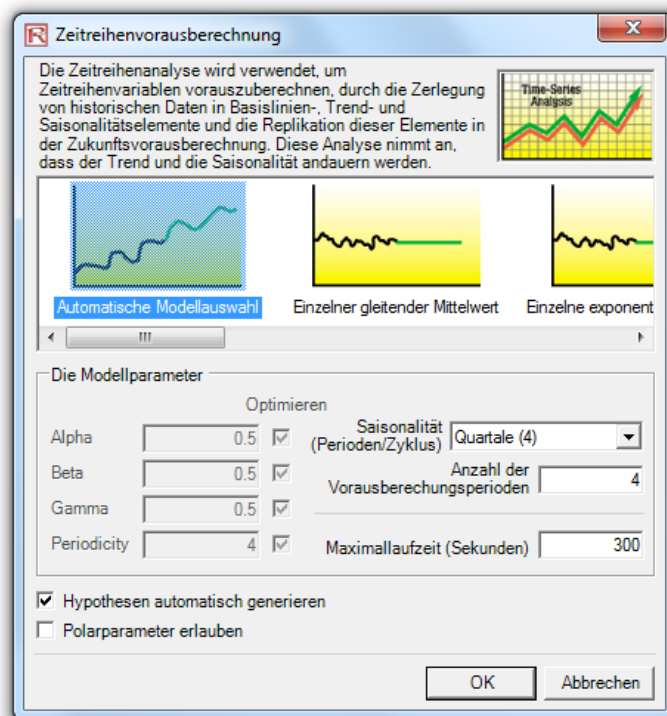


Bild 3.4—Zeitreihenanalyse

Interpretierung der Ergebnisse:

Bild 3.5 zeigt ein Beispiel von Ergebnissen, die unter Verwendung des Tools Vorausberechnung generiert wurden. Das verwendete Modell war das multiplikative Modell von Holt-Winters. Bitte bemerken Sie im Bild 3.5, dass sowohl die Modellanpassung als auch das Vorausberechnungsdiagramm darauf hindeuten, dass der Trend und die Saisonalität von dem multiplikativen Modell von Holt-Winters gut aufgenommen werden. Der Bericht der Zeitreihenanalyse stellt Ihnen Folgendes zur Verfügung: die relevanten optimierten Alpha-, Beta- und Gammaparameter, die Fehlermessungen, die angepassten Daten, die Vorausberechnungswerte und das Diagramm der angepassten Vorausberechnung. Die Parameter sind lediglich zur Einsichtnahme. Alpha erfasst den Memory-Effekt der Änderung des Basisniveaus im Laufe der Zeit; Beta ist der Trendparameter, der die Stärke des Trends misst; Gamma misst die Saisonalität der historischen Daten. Die Analyse zerlegt die historischen Daten in diese drei Elemente und stellt sie dann wieder zusammen, um die Zukunft vorherzusagen. Die angepassten Daten, unter Verwendung des wieder zusammengestellten Modells, stellen sowohl die historischen Daten als auch die angepassten Daten dar und zeigen wie nahe die Vorausberechnungen in der Vergangenheit sind (ein Verfahren genannt Zurückberechnung). Die Vorausberechnungswerte sind entweder Einzelpunktschätzungen oder -hypothesen (wenn die Option Hypothesen automatisch generieren ausgewählt ist und wenn ein Simulationsprofil vorhanden ist). Das Diagramm stellt diese historischen, angepassten und vorausberechneten Werte dar. Das Diagramm ist ein leistungsstarkes Kommunikations- und Anschauungstool, um die Güte des Vorausberechnungsmodells zu prüfen.

Anmerkungen:

Das Modul Zeitreihenanalyse enthält die 8 im Bild 3.3 angezeigten Zeitreihenmodelle. Sie können das spezifische auszuführende Modell, basierend auf den Kriterien Trend und Saisonalität, auswählen oder Auto-Modellauswahl aktivieren. Diese Option wird durch alle 8 Methoden automatisch iterieren, die Parameter optimieren und das bestpassende Modell für Ihre Daten finden. Alternativ, wenn Sie eins der 8 Modelle wählen, können Sie auch die Kontrollkästchen *Optimieren* deaktivieren und Ihre eigene Alpha-, Beta- und Gammaparameter eingeben. Lesen Sie *Modeling Risk, 2nd Edition: Applying Monte Carlo Simulation, Real Options Analysis, Forecasting, and Optimization* (Wiley, 2010) von Dr. Johnathan Mun für mehr Details über die technischen Spezifikationen dieser Parameter. Sie müssen außerdem die relevanten Saisonalitätsperioden eingeben, wenn Sie die automatische Modellauswahl oder irgendeines der Saisonalitätsmodelle auswählen. Der Saisonalitätsinput muss eine positive Ganzzahl sein (z.B., bei Quartalsdaten, geben Sie 4 als die Saisonanzahl oder Zyklen pro Jahr ein, oder bei Monatsdaten, geben Sie 12 ein). Als nächstes geben Sie die Anzahl der vorauszuberechnenden Perioden ein. Auch dieser Wert muss eine positive Ganzzahl sein. Die Maximallaufzeit ist auf 300 Sekunden eingestellt. Normalerweise sind keine Änderungen erforderlich. Allerdings, bei einer Vorausberechnung mit einer signifikanten Menge von historischen Daten, könnte die Analyse ein bisschen länger dauern und falls die Ausführungszeit diese Laufzeit überschreitet,

wir der Prozess beendet. Sie können auch bestimmen, dass die Vorausberechnung die Hypothesen automatisch generieren soll. Das heißt, anstelle von Einzelpunktschätzungen, werden die Vorausberechnungen Hypothesen sein. Als letztes, die Option Polarparameter erlaubt es Ihnen die Alpha-, Beta- und Gammaparameter so zu optimieren, dass sie Null und 1 einschließen. Einige Vorausberechnungssoftware erlauben diese Polarparameter, während andere das nicht tun. Risiko Simulator gibt Ihnen die Wahl welche zu benutzen. Normalerweise gibt es kein Grund, die Polarparameter zu verwenden.

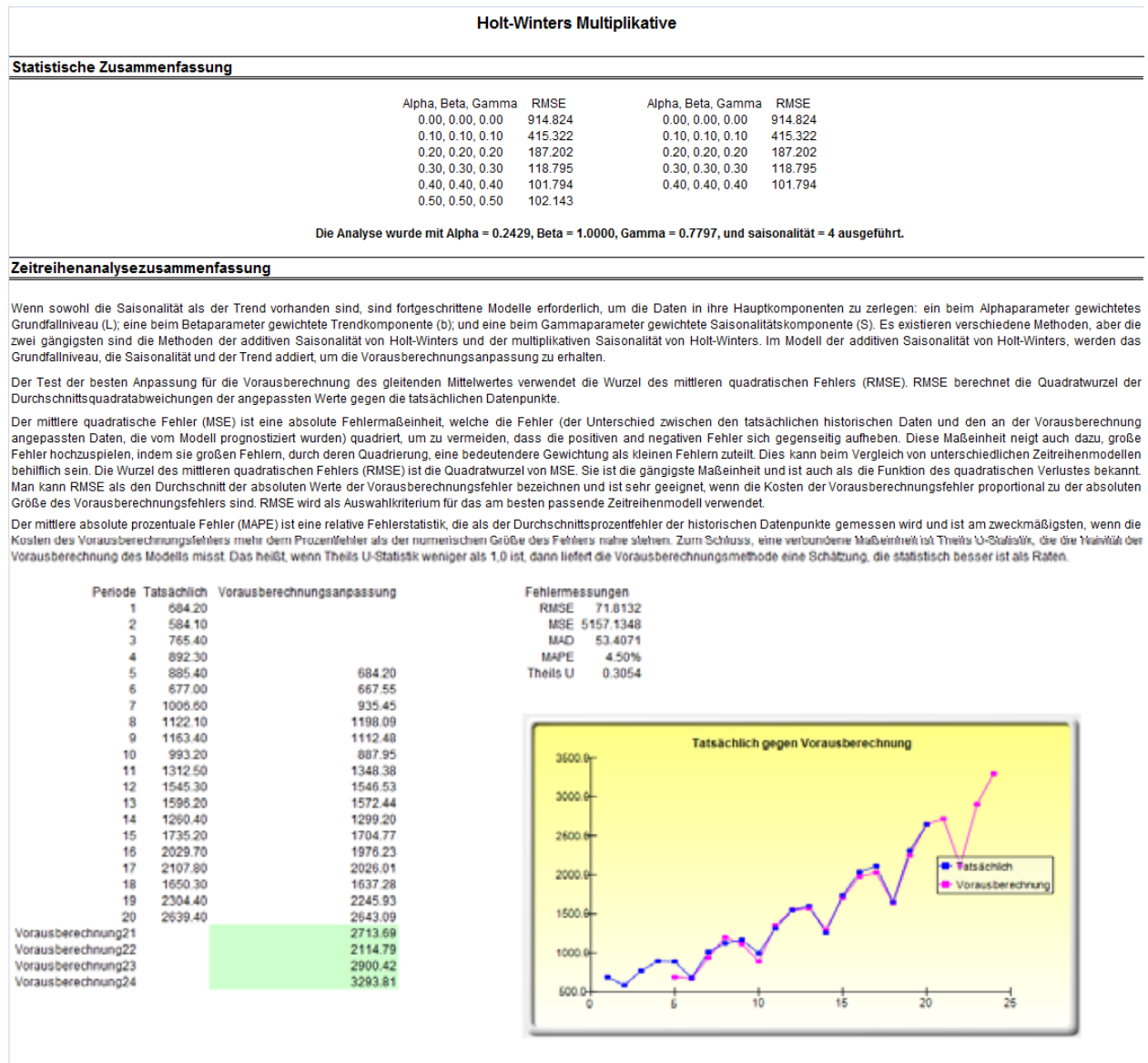


Bild 3.5—Beispiel - Bericht der Holt-Winters Vorausberechnung

Multivariate Regression

Theorie:

Es wird angenommen, dass der Benutzer genügend sachkundig in den Grundlagen der Regressionsanalyse ist. Die allgemeine bivariate lineare Regressionsgleichung nimmt die folgende Form an: $Y = \beta_0 + \beta_1 X + \varepsilon$, wobei β_0 der Achsenabschnitt, β_1 die Neigung und ε der Fehler-Term ist. Sie ist bivariate, da es nur zwei Variablen gibt, die Y oder abhängige Variable und die X oder unabhängige Variable, wobei X auch als der Regressor bekannt ist (eine bivariate Regression ist gelegentlich auch als eine univariate Regression bekannt, da es nur eine einzelne unabhängige Variable X gibt). Die abhängige Variable wird so genannt, weil sie von der unabhängigen Variablen *abhängt*. Verkaufseinnahmen, zum Beispiel, hängen vom Betrag der Marketingkosten ab, die für die Werbung und Förderung eines Produktes ausgegeben wurden. In diesem Beispiel sind die Verkäufe, die abhängige Variable und die Marketingkosten, die unabhängige Variable. Ein Beispiel einer bivariaten Regression kann folgendermaßen betrachtet werden: einfach die bestpassende Linie durch einen Satz von Datenpunkten in eine zweidimensionale Ebene einfügen, wie in der linken Seite des Bilds 3.6 angezeigt. In anderen Fällen kann man eine multivariate Regression ausführen, wobei es eine mehrfache oder n Anzahl von unabhängigen X Variablen gibt. Hier wird die allgemeine Regressionsgleichung die folgende Form annehmen: $Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 \dots + \beta_n X_n + \varepsilon$. In diesem Fall wird die bestpassende Linie innerhalb einer $n + 1$ dimensional Ebene sein.

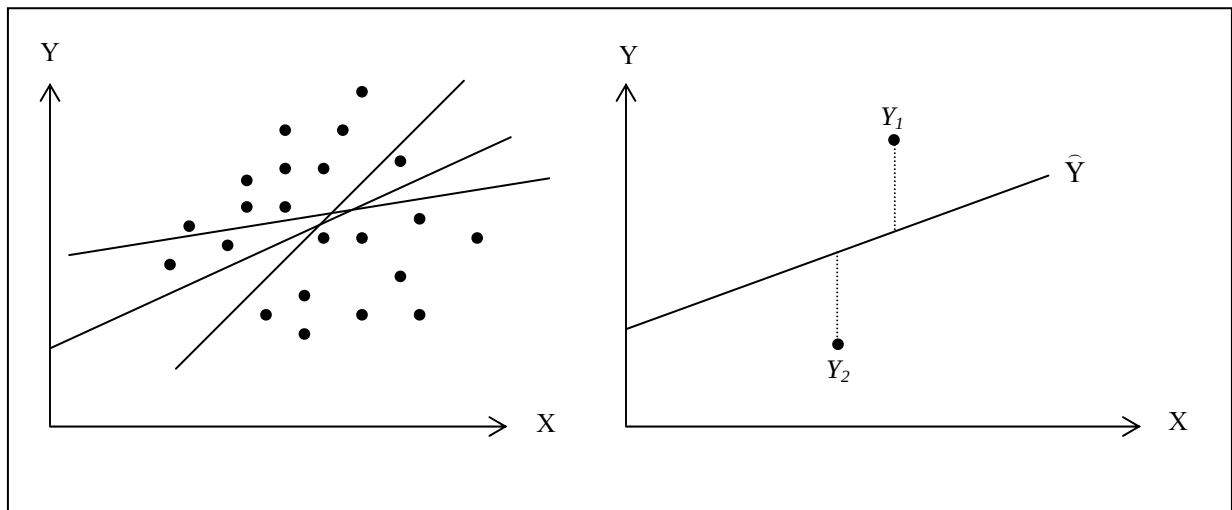


Bild 3.6—Bivariate Regression

Dennoch, die Anpassung einer Linie durch einen Satz von Datenpunkten in einem Streudiagramm wie im Bild 3.6 könnte in mehreren möglichen Linien resultieren. Die bestpassende Linie wird folgendermaßen definiert: Sie ist die einzelne einmalige Linie, welche die Gesamtvertikalfehler minimiert, sprich die Summe der absoluten Entfernungen zwischen den realen Datenpunkten (Y_i) und der geschätzten Linie ($\hat{Y}$) wie auf der rechten Seite des Bilds 3.6

angezeigt. Um die bestpassende Linie, welche die Fehler minimiert zu finden, ist eine anspruchsvollere Methode erforderlich: Diese ist die Regressionsanalyse. Die Regressionsanalyse findet also die einmalige bestpassende Linie entweder durch die Minimierung der Gesamtfehler oder durch die Berechnung von

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

wobei nur eine einmalige Linie, diese Summe der Quadratfehler minimiert. Die Fehler (Vertikalentfernung zwischen den realen Daten und der vorausberechneten Linie) werden quadriert, um zu verhindern, dass die negativen Fehler und die positiven Fehler sich aufheben. Die Lösung dieser Minimierungsaufgabe, mit Bezug auf die Neigung und den Achsabschnitt, erfordert erst die Berechnung von Ableitungen und ihre Einstellung auf Null:

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

was die Gleichungen der kleinsten Quadrate der bivariaten Regression ergibt:

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

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

Für die multivariate Regression wird die Analogie ausgeweitet, um die mehrfachen unabhängig Variablen zu berücksichtigen, wobei $Y_i = \beta_1 + \beta_2 X_{2,i} + \beta_3 X_{3,i} + \varepsilon_i$ und die geschätzten Neigungen können folgendermaßen berechnet werden:

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

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

Bei der Ausführung von multivariaten Regressionen, muss man sehr bei der Aufstellung und Interpretierung der Ergebnisse aufpassen. Zum Beispiel ist eine gute Kenntnis der ökonometrischen Modellierung erforderlich (z.B., die Identifizierung der Regressionsfallen wie etwa strukturelle Brüche, Multikollinearität, Heteroskedastizität, Autokorrelation, Spezifikationstests, Nichtlinearität und so weiter), bevor man ein geeignetes Modell aufbauen kann. Siehe *Modeling Risk, 2nd Edition: Applying Monte Carlo Simulation, Real Options Analysis, Forecasting, and Optimization* (Wiley, 2010) von Dr. Johnathan Mun für eine

ausführlichere Analyse und Diskussion von multivariaten Regressionen und wie man diese Regressionsfallen identifizieren kann.

Prozedur:

- Excel starten und nach Bedarf Ihre historischen Daten öffnen (das folgende Beispiel verwendet die **Mehrfache Regression** im Ordner Beispiele)
- Prüfen Sie, dass die Daten in Spalten angeordnet sind, wählen Sie erst den Gesamtdatenbereich einschließlich den Variablennamen aus und dann **Risiko Simulator | Vorausberechnung | Mehrfache Regression**
- Wählen Sie die abhängige Variable und prüfen Sie die relevanten Optionen (Verzögerungen, Schrittweise Regression, Nichtlineare Regression, und so weiter) und klicken Sie auf OK

Interpretierung der Ergebnisse:

Bild 3.8 zeigt das Beispiel eines Berichts der Ergebnisse einer multivariaten Regression. Der Bericht enthält alle Ergebnisse der Regression, der Varianzanalyse, des angepassten Diagramms und des Hypothesentests. Die technischen Details bezüglich der Interpretierung dieser Ergebnisse fallen außer des Bereichs dieses Handbuchs. Siehe *Modeling Risk, 2nd Edition: Applying Monte Carlo Simulation, Real Options Analysis, Forecasting, and Optimization* (Wiley 2010) von Dr. Johnathan Mun für eine ausführlichere Analyse und Diskussion von multivariaten Regressionen und die Interpretierung der Regressionsberichte.

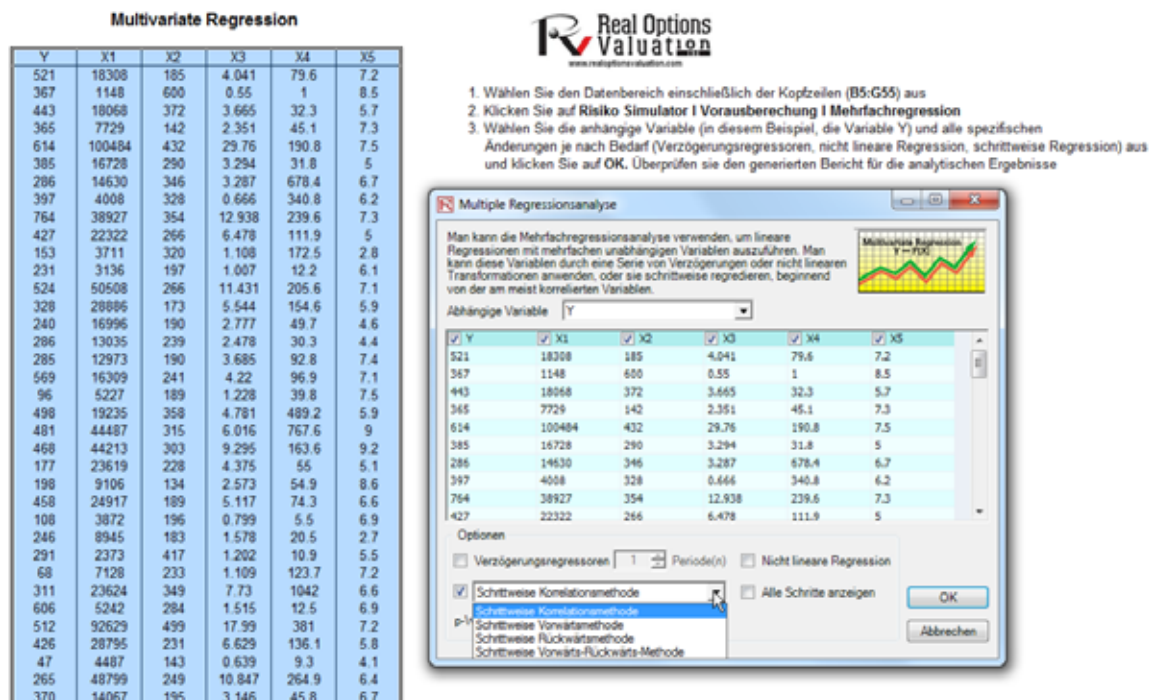


Bild 3.7—Ausführung einer multivariaten Regression

Regressionsanalysebericht

Regressionsstatistiken

R-Quadrat (Bestimmtheitskoeffizient)	0.3272
Korrigiertes R- Quadrat	0.2508
Mehrfaches R (Mehrfachkorrelationskoeffizient)	0.5720
Standardfehler der Schätzungen (SEy)	149.6720
Anzahl der Beobachtungen	50

Das R-Quadrat, oder der Bestimmtheitskoeffizient, zeigt an, dass 0.33 von der Variation in der abhängigen Variablen mittels der unabhängigen Variablen in dieser Regressionsanalyse erklärt und begründet werden kann. In einer Mehrfachregression jedoch, berücksichtigt das korrigierte R-Quadrat das Vorhandensein von zusätzlichen unabhängigen Variablen oder Regressoren und korrigiert diesen R-Quadrat-Wert zu einer genaueren Sicht der Regressionserklärungskraft. Daher kann nur 0.25 von der Variation in der abhängigen Variablen durch die Regressoren erklärt werden.

Der Mehrfachkorrelationskoeffizient (Multiple R) misst die Korrelation zwischen der tatsächlichen abhängigen Variablen (Y) und der geschätzten oder angepassten Variablen (Y) basierend auf der Regressionsgleichung. Dies ist auch die Quadratwurzel des Bestimmtheitskoeffizienten (R-Quadrat).

Der Standardfehler der Schätzungen (SEy) beschreibt die Dispersion der Datenpunkte über und unter der Regressionslinie oder -ebene. Dieser Wert wird als Teil der Berechnung verwendet, um später das Konfidenzintervall der Schätzungen zu erhalten.

Regressionsergebnisse

	Achsenabschnitt	X1	X2	X3	X4	X5
Koeffizienten	57.9555	-0.0035	0.4644	25.2377	-0.0086	16.5579
Standardfehler	108.7901	0.0035	0.2535	14.1172	0.1016	14.7996
t-Statistik	0.5327	-1.0066	1.8316	1.7877	-0.0843	1.1188
p-Wert	0.5969	0.3197	0.0738	0.0807	0.9332	0.2693
Untere 5%	-161.2966	-0.0106	-0.0466	-3.2137	-0.2132	-13.2687
Obere 95%	277.2076	0.0036	0.9753	53.6891	0.1961	46.3845

Freiheitsgrade

Freiheitsgrade für die Regression	5
Freiheitsgrade für das Residuum	44
Gesamtfreiheitsgrade	49

Hypothesentest

Kritische t-Statistik (99% Konfidenz mit Freiheitsgrade von 44)	2.6923
Kritische t-Statistik (95% Konfidenz mit Freiheitsgrade von 44)	2.0154
Kritische t-Statistik (90% Konfidenz mit Freiheitsgrade von 44)	1.6802

Die Koeffizienten liefern den geschätzten Achsenabschnitt und die geschätzten Steigungen der Regression. Zum Beispiel, die Koeffizienten sind Schätzungen der wahren b-Werte der Bevölkerung in der folgenden Regressionsgleichung $Y = b_0 + b_1X_1 + b_2X_2 + \dots + b_nX_n$. Der Standardfehler misst wie akkurat die vorausberechneten Koeffizienten sind, und die t-Statistiken sind die Verhältnisse von jedem vorausberechneten Koeffizienten zu seinem Standardfehler.

Die t-Statistik wird beim Hypothesentesten verwendet, wobei man die Nullhypothese (H_0) so einstellt, dass der reelle Mittelwert des Koeffizienten = 0, und die Alternativhypothese (H_a) so einstellt, dass der reelle Mittelwert des Koeffizienten nicht 0 gleicht. Ein t-Test wird ausgeführt und die berechnete t-Statistik wird mit den kritischen Werten an den relevanten Freiheitsgraden für das Residuum verglichen. Der t-Test ist sehr wichtig, da er berechnet, ob jeder der Koeffizienten statistisch signifikant in der Anwesenheit von anderen Regressoren ist. Das bedeutet, dass der t-Test statistisch nachprüft, ob man einen Regressor oder eine unabhängige Variable in der Regression behalten sollte oder nicht.

Der Koeffizient ist statistisch signifikant, wenn seine berechnete t-Statistik die kritische t-Statistik beim entsprechenden Freiheitsgrad (df) überschreitet. Die drei Hauptkonfidenzniveaus, die verwendet werden, um die Signifikanz zu testen, sind 90%, 95% und 99%. Wenn die t-Statistik eines Koeffizienten das kritische Niveau überschreitet, gilt er als statistisch signifikant. Alternativ, der p-Wert berechnet die Ereigniswahrscheinlichkeit jeder t-Statistik, was bedeutet, dass je kleiner der p-Wert, umso mehr signifikant der Koeffizient. Die üblichen Signifikanzniveaus für den p-Wert sind 0,01, 0,05 und 0,10, entsprechend den 99%, 95% und 99% Konfidenzniveaus.

Die Koeffizienten mit ihren p-Werten in blau hervorgehoben, zeigen an, dass sie statistisch signifikant beim 90% Konfidenz- oder 0,10 Alphaniveau sind, während die in rot hervorgehobenen anzeigen, dass sie nicht statistisch signifikant bei irgendwelchen anderen Alphaniveaus sind.

Varianzanalyse

	Summen der Quadrate	Mittelwert der Quadrate	F-Statistik	p-Wert	Hypothesentest	
Regression	479388.49	95877.70	4.28	0.0029	Kritische F-Statistik (99% Konfidenz mit Freiheitsgrade von 5	3.4651
Residuum	985675.19	22401.71			Kritische F-Statistik (95% Konfidenz mit Freiheitsgrade von 5	2.4270
Summe	1465063.68				Kritische F-Statistik (90% Konfidenz mit Freiheitsgrade von 5	1.9828

Die Tabelle der Varianzanalyse (ANOVA) liefert einen F-Test der gesamten statistischen Signifikanz des Regressionsmodells. Anstatt die individuellen Regressoren, wie beim t-Test, zu betrachten, prüft der F-Test alle die statistischen Eigenschaften der geschätzten Koeffizienten. Die F-Statistik wird als das Verhältnis des Mittelwertes der Quadrate der Regression zu dem Mittelwert der Quadrate des Residuums berechnet. Der Zähler misst wie viel von der Regression erklärt wird, während der Nenner den unerklärten Teil misst. Deshalb, je größer die F-Statistik, umso mehr signifikant das Modell. Der entsprechende p-Wert wird berechnet, um die Nullhypothese (H_0), wo alle Koeffizienten gleichzeitig Null gleichen, gegen die Alternativhypothese (H_a), wo sie alle gleichzeitig abweichend von Null sind, zu testen, was auf ein signifikantes Gesamtregressionsmodell hinweist. Wenn der p-Wert kleiner als die 0,01, 0,05 oder 0,10 Alphasignifikanz ist, dann ist die Regression signifikant. Man kann dieselbe Methode auf die F-Statistik anwenden, indem man die berechnete F-Statistik mit den kritischen F-Werten an verschiedenen Signifikanzniveaus vergleicht.

Vorausberechnung

Periode	sächlich (Y)	erechnung (F)	Fehler (E)
1	521.0000	299.5124	221.4876
2	367.0000	487.1243	(120.1243)
3	443.0000	353.2789	89.7211
4	365.0000	276.3296	88.6704
5	614.0000	776.1336	(162.1336)
6	385.0000	298.9993	86.0007
7	286.0000	354.8718	(68.8718)
8	397.0000	312.6155	84.3845
9	764.0000	529.7550	234.2450
10	427.0000	347.7034	79.2966
11	153.0000	266.2526	(113.2526)
12	231.0000	264.6375	(33.6375)
13	524.0000	406.8009	117.1991
14	328.0000	272.2226	55.7774
15	240.0000	231.7882	8.2118
16	286.0000	257.8862	28.1138
17	285.0000	314.9521	(29.9521)
18	569.0000	335.3140	233.6860
19	96.0000	282.0356	(186.0356)
20	498.0000	370.2062	127.7938
21	481.0000	340.8742	140.1258
22	468.0000	427.5118	40.4882

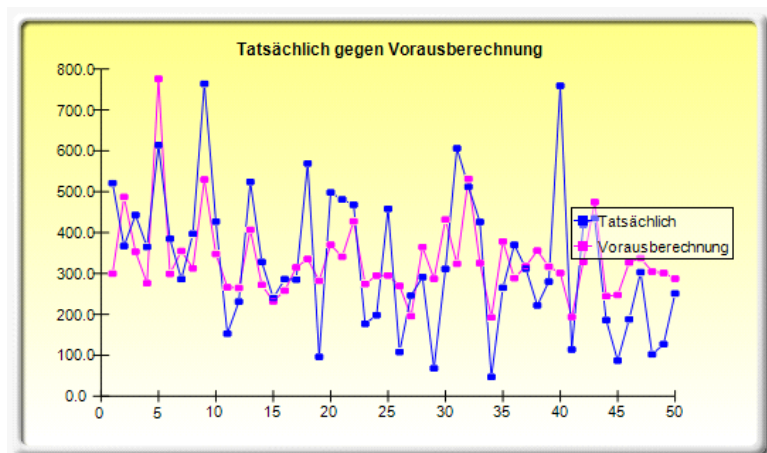


Bild 3.8—Ergebnisse der multivariaten Regression

Stochastische Vorausberechnung

Theorie:

Ein stochastischer Prozess ist nichts mehr als eine mathematisch definierte Gleichung, die eine Reihe von Ausgängen im Laufe der Zeit kreieren kann, Ausgänge die nicht deterministisch sind. Das heißt, eine Gleichung oder ein Prozess, die/der nicht irgendeiner einfachen erkennbaren Regel wie etwa „der Preis steigt X Prozent jedes Jahr“ oder „Einnahmen steigen um diesen Faktor von X plus Y Prozent“ folgen. Ein stochastischer Prozess ist definitionsgemäß nichtdeterministisch. Man kann Zahlen in die Gleichung eines stochastischen Prozesses eingeben und unterschiedliche Ergebnisse jedes Mal erhalten. Zum Beispiel, der Pfad eines Aktienpreises ist stochastisch und man kann nicht zuverlässig den Aktienpreispfad mit irgendeiner Gewissheit vorausberechnen. Allerdings ist die Preisevolution im Laufe der Zeit in einem Prozess eingewickelt, der diese Preise generiert. Der Prozess ist festgelegt und vorherbestimmt, aber die

Ausgänge sind es nicht. Unter Verwendung einer stochastischen Simulation können wir deshalb mehrfache Preispfade kreieren, eine statistische Stichprobenentnahme dieser Simulationen erhalten und Schlussfolgerungen über die potentiellen Pfade, die der reale Preis nehmen könnte, ziehen, gegeben die Natur und Parameter des stochastischen Prozesses, der zur Generierung der Zeitreihen verwendet wurde. Drei stochastische Grundprozesse sind in *Vorausberechnungs-Tool* von *Risiko Simulator* enthalten, einschließlich der geometrischen Brownsche Bewegung oder Irrfahrt, welche, auf Grund ihrer Einfachheit und umfangreichen Anwendungen, die gängigste und am häufigsten verwendete Methode ist. Die anderen zwei stochastischen Prozesse sind Rückkehr zum Mittelwert und Sprung-Diffusion .

Das Interessante bei der Simulation mit einem stochastischen Prozess ist, dass historische Daten nicht unbedingt erforderlich sind. Das heißt, das Modell muss sich nicht an irgendwelchen Sätzen von historischen Daten anpassen. Man berechnet einfach die Erwartungserträge und die Volatilität der historischen Daten oder man schätzt sie unter Verwendung von vergleichbaren externen Daten oder man macht Hypothesen über diese Werte. Siehe *Modeling Risk: Applying Monte Carlo Simulation, Real Options Analysis, Forecasting, and Optimization*, 2nd Edition (Wiley 2010) von Dr. Johnathan Mun für mehr Details über wie diese Inputs berechnet werden (z.B., Rückkehr-zum-Mittelwert Rate, Sprungwahrscheinlichkeiten, Volatilität und so weiter).

Prozedur:

- Starten Sie das Modul mittels ***Risiko Simulator | Vorausberechnung | Stochastische Prozesse***
- Wählen Sie den gewünschten Prozess, geben Sie die erforderlichen Inputs ein und klicken Sie einige Male auf Diagramm aktualisieren, um sicher zu gehen, dass der Prozess sich so verhält, wie Sie es erwarten. Schließlich klicken Sie auf OK (Bild 3.9)

Interpretierung der Ergebnisse:

Bild 3.10 zeigt die Ergebnisse des Beispiels eines stochastischen Prozesses. Das Diagramm zeigt einen Beispielsatz von Iterationen, während der Bericht die Grundlagen der stochastischen Prozesse erklärt. Außerdem werden die Vorausberechnungswerte (Mittelwert und Standardabweichung) für jede Zeitperiode bereitgestellt. Unter Verwendung dieser Werte können Sie entscheiden, welche Zeitperiode für Ihre Analyse relevant ist und Hypothesen, unter Verwendung der Normalverteilung, auf Basis dieser Werte des Mittelwerts und der Standardabweichung aufstellen. Sie können dann diese Hypothesen in Ihrem eigenen angepassten Modell simulieren.

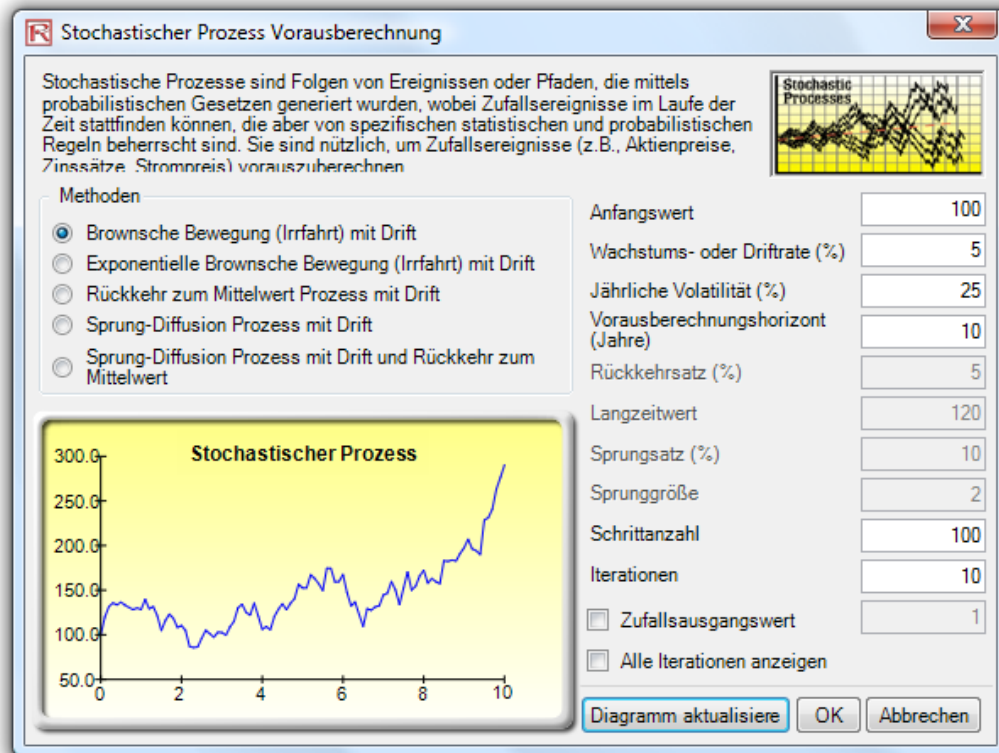


Bild 3.9 –Stochastischer Prozess Vorausberechnung

Stochastischer Prozess Vorausberechnung

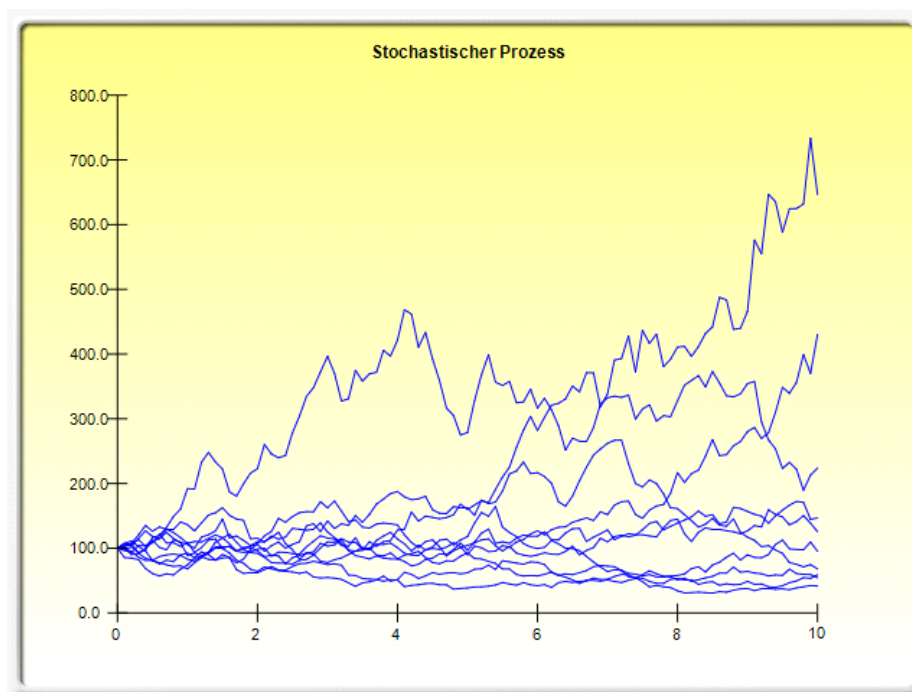
Statistische Zusammenfassung

Ein stochastischer Prozess ist eine Folge von Ereignissen oder Pfaden, die mittels probabilistischen Gesetzen generiert wurden. Das heißt, Zufallsereignisse können im Laufe der Zeit stattfinden, aber sie werden von spezifischen statistischen und probabilistischen Regeln beherrscht. Die wichtigsten stochastischen Prozesse sind, unter anderem, die Irrfahrt oder Brownsche Bewegung, die Rückkehr zum Mittelwert und die Sprung-Diffusion. Man kann diese Prozesse verwenden, um eine Vielzahl von Variablen, die scheinbar Zufallstrends folgen aber doch durch probabilistische Gesetze eingeschränkt sind, vorzuberechnen.

Der Prozess mit Irrfahrt Brownsche Bewegung kann verwendet werden, um Aktienpreise, Warenpreise und andere stochastische Zeitreihendaten vorzuberechnen, gegeben eine Drift oder Wachstumsrate und eine Volatilität um den Driftpfad. Der Prozess mit Rückkehr zum Mittelwert kann verwendet werden, um die Fluktuationen des Irrfahrtprozesses zu reduzieren, indem es dem Pfad erlaubt einen Langzeitwert anzupfeilen. Dies macht ihn nützlich, um Zeitreihenvariablen vorzuberechnen, die eine Langzeitrage haben, sowie Zinssätze und Inflationsraten (dies sind Langzeitzielraten der Regulierungsbehörden oder des Markts). Der Prozess mit Sprung-Diffusion ist nützlich, um Zeitreihendaten vorzuberechnen, wenn die Variable gelegentlich Zufallssprünge aufweisen kann, sowie Erdöl- oder Strompreise (diskrete exogene Ereignisshocks können Preise nach oben oder nach unten springen lassen). Letztlich, man kann diese drei stochastischen Prozesse nach Wunsch mischen und anpassen.

Die Ergebnisse auf der rechten Seite zeigen die Durchschnitts- und Standardabweichung von allen Iterationen, die bei jedem Zeitschritt generiert werden. Wenn die Option 'Alle Iterationen anzeigen' ausgewählt ist, wird jeder Iterationspfad in einem separaten Arbeitsblatt angezeigt. Das nachstehende generierte Diagramm zeigt einen Stichprobensatz der Iterationspfade.

Stochastischer Prozess: Brownsche Bewegung (Irrfahrt) mit Drift					
Anfangswert	100	Schritte	100.00	Sprungsatz	N/A
Driftrate	5.00%	Iterationen	10.00	Sprunggröße	N/A
Volatilität	25.00%	Rückkehrrate	N/A	Zufallsausgangswert	1306519657
Horizont	10	Langzeitwert	N/A		



Zeit	Mittelwert	StAbw
0.0000	100.00	0.00
0.1000	100.87	6.57
0.2000	100.12	8.79
0.3000	100.08	12.94
0.4000	99.50	19.91
0.5000	99.65	21.68
0.6000	100.20	24.69
0.7000	102.62	25.81
0.8000	103.90	28.17
0.9000	105.67	28.74
1.0000	105.66	36.21
1.1000	104.25	33.97
1.2000	114.19	45.00
1.3000	120.00	49.35
1.4000	122.41	44.46
1.5000	123.85	42.21
1.6000	115.11	31.76
1.7000	106.51	34.22
1.8000	106.82	39.17
1.9000	105.31	42.12
2.0000	107.41	44.44
2.1000	110.51	55.22
2.2000	110.45	51.60
2.3000	112.45	51.69
2.4000	108.40	53.07
2.5000	114.89	63.29
2.6000	117.93	71.03
2.7000	126.30	78.55
2.8000	129.16	82.47
2.9000	134.77	90.00
3.0000	136.13	96.99
3.1000	133.02	89.44
3.2000	127.19	76.48
3.3000	120.84	79.36
3.4000	127.21	93.30
3.5000	121.15	88.24
3.6000	124.33	91.76
3.7000	127.81	92.72
3.8000	133.38	102.72
3.9000	133.97	100.65
4.0000	135.69	108.45
4.1000	136.85	122.63
4.2000	135.26	121.58
4.3000	128.94	106.34
4.4000	133.88	112.62
4.5000	128.41	99.22
4.6000	124.18	89.40
4.7000	119.63	77.14
4.8000	120.70	74.95
4.9000	121.51	66.94
5.0000	120.28	67.51
5.1000	129.61	79.29
5.2000	139.91	91.56
5.3000	142.06	99.69

Bild 3.10—Ergebnisse der stochastischen Vorausberechnung

Nichtlineare Extrapolation

Theorie:

Bei der Extrapolation geht es, um die Ausführung von statistischen Projektionen, indem man historische Trends verwendet, welche für einen spezifischen Zeitraum in die Zukunft projiziert werden. Sie wird nur für Zeitreihenvorausberechnungen verwendet. Für Querschnitts- oder Mischpaneldaten (Zeitreihen mit Querschnittsdaten), ist eine multivariate Regression geeigneter. Diese Methodologie ist nützlich, wenn wesentliche Änderungen nicht erwartet werden, das heißt, wenn man erwartet, dass die Kausalfaktoren konstant bleiben werden oder wenn die Kausalfaktoren einer Situation nicht klar verstanden sind. Sie hilft auch die Einführung von persönlichen Verzerrungen (Bias) im Prozess zu verhindern. Die Extrapolation ist ziemlich zuverlässig, relativ einfach und preiswert. Allerdings produziert die Extrapolation, die annimmt, dass die jüngsten und historischen Trends andauern werden, große Vorausberechnungsfehler, wenn sich Unterbrechungen während des projizierten Zeitraums ereignen. In anderen Worten, die reine Extrapolation von Zeitreihen nimmt an, dass alles was wir wissen müssen, in den historischen Werten der zu vorausberechnenden Zeitreihen enthalten ist. Wenn wir annehmen, dass das Verhalten in der Vergangenheit ein guter Prädiktor von zukünftigem Verhalten ist, ist die Extrapolation interessant. Dies macht sie zu einer nützlichen Methode, wenn nur viele Kurzzeitvorausberechnungen benötigt werden.

Diese Methodologie schätzt die $f(x)$ Funktion für alle beliebigen x Werte, indem sie eine glatte nichtlineare Kurve durch alle x Werte interpoliert und, unter Verwendung dieser glatten Kurve, zukünftige x Werte über den historischen Datensatz hinaus extrapoliert. Die Methodologie verwendet entweder die polynomiale Funktionsform oder die rationale Funktionsform (ein Verhältnis von zwei Polynomen). Normalerweise ist eine polynomiale Funktionsform ausreichend für Daten mit guten Verhalten, aber rationale Funktionsformen sind gelegentlich akkurater (insbesondere mit Polarfunktionen, das heißt, Funktionen mit Nenner, welche sich der Null annähern).

Prozedur:

- ❏ Excel starten und bei Bedarf Ihre historischen Daten öffnen (das nachstehend angezeigte Beispiel verwendet die Datei ***Nicht lineare Extrapolation*** aus dem Ordner Beispiele)
- ❏ Wählen Sie erst die Zeitreihendaten und dann ***Risiko Simulator | Vorausberechnung | Nicht lineare Extrapolation***
- ❏ Wählen Sie den Extrapolationstyp (automatische Auswahl, Polynomfunktion oder Rationalfunktion), geben Sie die Anzahl der gewünschten Vorausberechnungsperioden ein (Bild 3.11) und klicken Sie auf OK

Interpretierung der Ergebnisse:

Der im Bild 3.12 angezeigter Bericht der Ergebnisse präsentiert die extrapolierten Vorausberechnungswerte, die Fehlermessungen und die graphische Darstellung der

Extrapolationsergebnisse. Man sollte die Fehlermessungen verwenden, um die Gültigkeit der Vorausberechnung zu prüfen. Dies ist besonders wichtig, wenn sie verwendet werden, um die Vorausberechnungsqualität und die Extrapolationsgenauigkeit mit der Zeitreihenanalyse zu vergleichen.

Bemerkungen:

Wenn die historischen Daten glatt sind und irgendwelchen nichtlinearen Muster und Kurven folgen, ist die Extrapolation besser als die Zeitreihenanalyse. Wenn die Datenmuster jedoch saisonabhängige Zyklen und einem Trend folgen, liefert die Zeitreihenanalyse bessere Ergebnisse

Nicht lineare Extrapolation

Bitte beachten Sie, dass es bei der nicht linearen Extrapolation, um die Ausführung von statistischen Projektionen geht. Es werden historische Trends verwendet, welche für einen Spezifischen Zeitraum in die Zukunft projiziert werden. Sie wird nur bei der Vorausberechnung von Zeitreihen verwendet. Die Extrapolation ist ziemlich zuverlässig, relativ einfach und kostengünstig. Da die Extrapolation allerdings annimmt, dass die jüngsten und die historischen Trends andauern werden, produziert sie große Vorausberechnungsfehler, wenn sich Unstetigkeiten während des projizierten Zeitraums ereignen.

Historische Verkaufseinnahmen Polynom Wachstumsraten

Jahr	Monat	Periode	Umsatz
2004	1	1	\$1.00
2004	2	2	\$6.73
2004	3	3	\$20.52
2004	4	4	\$45.25
2004	5	5	\$83.59
2004	6	6	\$138.01
2004	7	7	\$210.87
2004	8	8	\$304.44
2004	9	9	\$420.89
2004	10	10	\$562.34
2004	11	11	\$730.85
2004	12	12	\$928.43

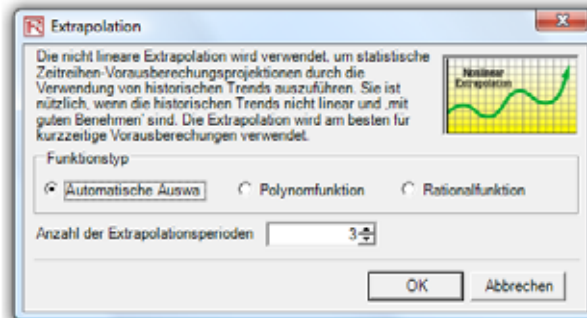


Bild 3.11 –Eine nicht lineare Extrapolation ausführen

Nicht lineare Extrapolation

Statistische Zusammenfassung

Die Extrapolation umfasst die Ausführung von statistischen Vorausberechnungen unter Verwendung von historischen Trends, die für einen angegebenen Zeitraum in die Zukunft projiziert sind. Sie wird nur für Zeitreihenvorausberechnungen verwendet. Für Querschnitts- oder Mischpanelndaten (Zeitreihen mit Querschnittsdaten) ist die multivariate Regression geeigneter. Diese Methodik ist nützlich, wenn keine bedeutenden Änderungen erwartet werden, das heißt, wenn man vermutet, dass Kausalfaktoren konstant bleiben werden oder wenn die Kausalfaktoren einer Situation nicht deutlich zu verstehen sind. Sie hilft auch die Einführung von persönlichen Verzerrungen in den Prozess zu verhindern. Die Extrapolation ist ziemlich zuverlässig, relativ einfach und preiswert. Allerdings produziert die Extrapolation, die annimmt, dass jüngste und historische Trends fortauern werden, große Vorausberechnungsfehler, wenn sich Unterbrechungen während des projizierten Zeitraums ereignen. Das heißt, die reine Extrapolation einer Zeitreihe nimmt an, dass alles was wir wissen müssen sich in den historischen Werten der zu vorausberechnenden Reihe befindet. Wenn wir annehmen, dass das vergangene Verhalten ein guter Prädiktor von zukünftigem Verhalten ist, ist die Extrapolation ansprechend. Das macht sie zu einer nützlichen Methode, wenn nur viele kurzzeitige Vorausberechnungen benötigt werden.

Diese Methodik schätzt die $f(x)$ Funktion für irgendeinen beliebigen x -Wert, indem sie eine glatte nicht lineare Kurve durch alle x -Werte interpoliert. Dann verwendet sie diese glatte Kurve, um zukünftige x -Werte jenseits des historischen Datensatzes zu extrapolieren. Die Methodik verwendet entweder die Polynomfunktionsform oder die Rationalfunktionsform (ein Verhältnis von zwei Polynomen). Typisch ist eine Polynomfunktionsform ausreichend für Daten 'mit gutem Benehmen', allerdings sind Rationalfunktionsformen manchmal akkurater (insbesondere mit Polarfunktionen, sprich Funktionen mit Nennern, die sich der Null annähern).

Periode	Tatsächlich	Vorausberechnungsanpassung
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
Vorausberechnung (Prognose) 13		1157.03

Fehlerrmessungen	
RMSE	19.6799
MSE	387.2974
MAD	10.2095
MAPE	31.56%
Theils U	1.1210

Funktionstyp: Rationale

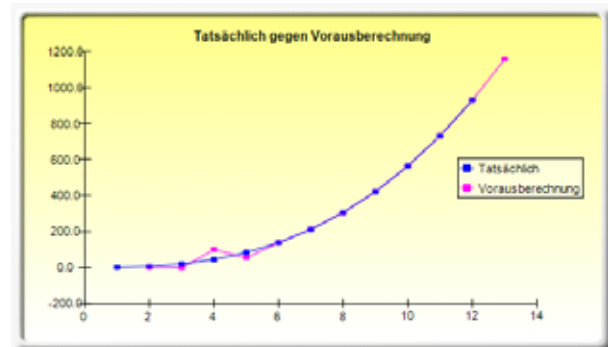


Bild 3.12—Ergebnisse einer nicht linearen Extrapolation

Box-Jenkins ARIMA Fortgeschrittene Zeitreihen

Theorie:

Ein sehr leistungsstarkes fortgeschrittenes Zeitreihen-Vorausberechnungstool ist das Verfahren ARIMA oder *autoregressiver integrierter gleitender Mittelwert*. Die ARIMA Vorausberechnung sammelt drei einzelne Tools in ein umfassendes Modell. Das erste Toolsegment ist der Begriff Autoregression oder „AR“. Dieser Begriff entspricht der Anzahl der verzögerten Werte des Residuums in das unbedingte Vorausberechnungsmodell. Im Wesentlichen, das Modell erfasst die historische Variation der tatsächlichen Daten verglichen mit einem Vorausberechnungsmodell und verwendet diese Variation oder dieses Residuum, um ein besseres vorausberechnendes Modell zu kreieren. Das zweite Toolsegment ist der Begriff Integrationsbefehl oder „I“. Dieser Integrationsbegriff entspricht der Anzahl der Differenzierungen, welche die zu vorausberechnenden Zeitreihen durchgehen müssen. Dieses Element berücksichtigt alle nicht linearen Wachstumsraten, die in den Daten vorhanden sind. Das dritte Toolsegment ist der Begriff gleitender Mittelwert oder „MA“. Dieser ist sozusagen der gleitende Mittelwert der verzögerten Vorausberechnungsfehler. Durch die Inkorporierung dieser verzögerten Vorausberechnungsfehler, lernt das Modell im Grunde von seinen Vorausberechnungsfehlern oder -irrtümern und korrigiert Sie durch eine „gleitender Mittelwert“ Berechnung. Das ARIMA Modell folgt die Box-Jenkins Methodologie: jeder Begriff repräsentiert die beim Modellaufbau

durchgeführten Schritte, bis nur Zufallsrauschen übrig geblieben ist. Die ARIMA Modellierung verwendet außerdem Korrelationsverfahren bei der Generierung von Vorausberechnungen. Man kann ARIMA verwenden, um Muster zu modellieren, die in gezeichneten Daten nicht sichtbar sein könnten. Außerdem kann man ARIMA Modelle mit exogenen Variablen mischen, aber achten Sie darauf, dass die exogenen Variablen genügend Datenpunkte haben, um die zusätzliche Anzahl von vorauszuberechnenden Perioden zu decken. Zum Schluss, seien Sie sich bewusst, dass auf Grund der Komplexität der Modelle, die Ausführung dieses Moduls eine längere Zeit brauchen könnte.

Es gibt viele Gründe, warum ein ARIMA Modell den einfachen Zeitreihenanalysen und multivariaten Regressionen überlegen ist. Die gemeinschaftliche Feststellung bei Zeitreihenanalysen und multivariaten Regressionen ist, dass die Restfehler mit ihren eigenen verzögerten Werten korreliert werden. Diese serielle Korrelation verletzt die Standardhypothese der Regressionstheorie, dass Störungen nicht mit anderen Störungen korreliert sind. Die mit der seriellen Korrelation verbundenen Hauptprobleme sind:

- Regressionsanalysen und elementare Zeitreihenanalysen sind nicht länger effizient unter den verschiedenen linearen Schätzern. Allerdings, da die Restfehler bei der Vorausberechnung der Restfehler helfen können, kann man diese Informationen benutzen, um unter Verwendung von ARIMA eine bessere Vorausberechnung der abhängigen Variablen zu gestalten.
- Standardfehler, die unter Verwendung der Regressions- und Zeitreihenformeln berechnet werden, sind nicht korrekt und normalerweise untertrieben. Wenn außerdem verzögerte abhängige Variablen als die Regressoren eingestellt sind, sind die Regressionsschätzungen verzerrt und uneinheitlich, können aber unter Verwendung von ARIMA korrigiert werden.

Autoregressiver integrierter gleitender Mittelwert oder ARIMA(p,d,q) Modelle sind die Erweiterung des AR Modells, die drei Komponenten verwenden, um die serielle Korrelation in den Zeitreihendaten zu modellieren. Die erste Komponente ist der Begriff Autoregression (AR). Das AR(p) Modell verwendet die p-Verzögerungen der Zeitreihen in der Gleichung. Ein AR(p) Modell trägt die Form: $y_t = a_1 y_{t-1} + \dots + a_p y_{t-p} + e_t$. Die zweite Komponente ist der Begriff Integrationsbefehl (d). Jeder Integrationsbefehl entspricht einer Differenzierung der Zeitreihen. I(1) bedeutet, die Daten einmal zu differenzieren. I(d) bedeutet, die Daten d Male zu differenzieren. Die dritte Komponente ist der Begriff gleitender Mittelwert (MA). Das MA(q) Modell verwendet die q-Verzögerungen der Vorausberechnungsfehler, um die Vorausberechnung zu verbessern. Ein MA(q) Modell trägt die Form: $y_t = e_t + b_1 e_{t-1} + \dots + b_q e_{t-q}$. Zum Schluss, ein ARMA(p,q) Modell trägt die kombinierte Form: $y_t = a_1 y_{t-1} + \dots + a_p y_{t-p} + e_t + b_1 e_{t-1} + \dots + b_q e_{t-q}$.

Prozedur:

- Excel starten und Ihre Daten eingeben oder ein existierendes vorauszuberechnendes Arbeitsblatt mit historischen Daten öffnen (das nachstehend angezeigte Beispiel verwendet die Beispielsdatei **Zeitreihen ARIMA**)
- Wählen Sie die Zeitreihendaten und dann **Risiko Simulator | Vorausberechnung | ARIMA**
- Geben Sie die relevanten P, D, and Q Parameter (nur positive Ganzzahlen) und die Anzahl der gewünschten Vorausberechnungsperioden ein und klicken Sie auf **OK**

Bemerkung für ARIMA und AUTO-ARIMA:

- Für ARIMA und Auto-ARIMA: Sie können zukünftige Perioden modellieren und berechnen entweder unter Verwendung lediglich der abhängigen Variablen (Y), das heißt, die **Zeitreihenvariable** alleine, oder Sie können zusätzliche exogene Variablen ($X_1, X_2, \dots, X_n$) hinzufügen, genau wie in einer Regressionsanalyse wo man mehrfache unabhängige Variablen hat. Wenn Sie nur die Zeitreihenvariable (Y) verwenden, können Sie beliebig viele Vorausberechnungsperioden ausführen. Wenn Sie jedoch exogene Variablen (X) hinzufügen, bemerken Sie bitte, dass die Vorausberechnungsperioden auf die Anzahl der Datenperioden der exogenen Variablen minus den Datenperioden der Zeitreihenvariablen beschränkt sind. Zum Beispiel, Sie können nur bis zu 5 Perioden vorausberechnen, wenn Sie historische Zeitreihendaten von 100 Perioden und nur wenn Sie exogene Variablen von 105 Perioden haben (100 historische Perioden zur Angleichung der Zeitreihenvariablen und 5 zusätzliche Zukunftsperioden von unabhängigen exogenen Variablen, um die abhängige Zeitreihenvariable vorauszuberechnen).

Interpretierung der Ergebnisse:

Bei der Interpretierung der Ergebnisse eines ARIMA Modells, sind die meisten Spezifikationen identisch mit der multivariaten Regressionsanalyse (siehe *Modeling Risk*, 2nd Edition by Dr. Johnathan Mun für mehr technische Details über die Interpretierung der multivariaten Regressionsanalyse und ARIMA Modelle). Allerdings gibt es einige zusätzliche Sätze von Ergebnissen, welche für die ARIMA Analyse spezifisch sind, wie im Bild 3.14 angezeigt. Der erste Satz ist das Hinzufügen des Akaike-Informationskriterium (AIC) und des Schwarz-Kriteriums (SC), die oft in der ARIMA Modellauswahl und -identifizierung verwendet werden. Das heißt, AIC und SC werden verwendet, um festzustellen, ob ein bestimmtes Modell mit einem spezifischen Satz von p, d und q Parametern eine gute statistische Anpassung ist. SC verhängt eine größere Strafe für zusätzliche Koeffizienten als AIC, aber generell sollte man das Modell mit den niedrigsten Werten von AIC und SC wählen. Schließlich wird ein zusätzlicher Satz von Ergebnissen, welche die Autokorrelation (AC) und partielle Autokorrelation (PAC) Statistiken genannt werden, im ARIMA Bericht zur Verfügung gestellt.

Zum Beispiel, wenn die Autokorrelation AC(1) ungleich Null ist, bedeutet das, dass die Reihe in der 1. Ordnung seriell korreliert ist. Wenn AC mehr oder weniger geometrisch mit zunehmender Verzögerung wegbricht, bedeutet das, dass die Reihe einem autoregressiven Prozess der unteren

Ordnung folgt. Wenn AC nach einer kleinen Anzahl von Verzögerungen auf Null fällt, bedeutet das, dass die Reihe einem gleitenden Mittelwert Prozess der unteren Ordnung folgt. PAC dagegen misst die Korrelation der Werte, die k Perioden auseinander sind nach Entfernung der Korrelation von den Zwischenverzögerungen. Wenn die Struktur der Autokorrelation von einer Autoregression mit einer Ordnung kleiner als k erfasst werden kann, dann wird die partielle Autokorrelation bei der Verzögerung k nahe Null liegen. Die Ljung-Box Q-Statistiken und deren p-Werte bei der Verzögerung k werden auch bereitgestellt, wobei die getestete Nullhypothese so ist, dass es keine Autokorrelation bis hin zur Ordnung k gibt. Die punktierten Linien in den Diagrammen der Autokorrelationen sind die ungefähren zwei Standardfehlergrenzen. Wenn sich die Autokorrelation innerhalb diesen Grenzen befindet, ist sie nicht signifikant abweichend von Null bei ungefähr dem 5% Signifikanzniveau. Das richtige ARIMA Modell zu finden erfordert Übung und Erfahrung. AC, PAC, SC und AIC sind sehr nützliche Diagnosetools, die behilflich bei der Identifizierung der richtigen Modellspezifikation sind.

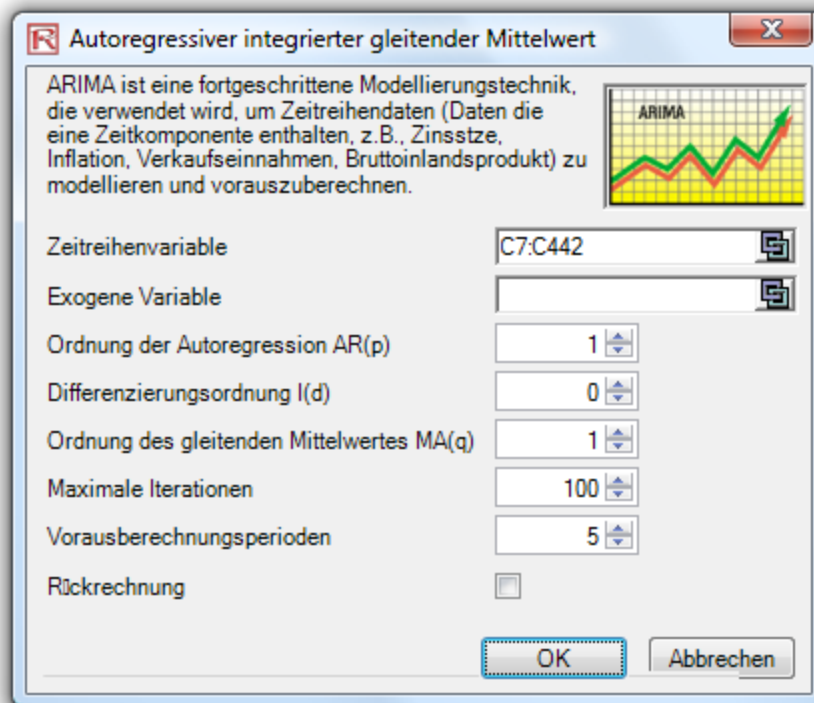


Bild 3.13A - Box Jenkins ARIMA Vorausberechnungstool

ARIMA (Autoregressives integriertes gleitendes Mittel)

Regressionsstatistiken			
R-Quadrat (Bestimmtheitskoeffizient)	0.9999	Akaike-Information-Kriterium (AIC)	4.6213
Korrigiertes R-Quadrat	0.9999	Schwarz-Kriterium (SC)	4.6632
Mehrfaches R (Mehrfachkorrelationskoeffizient)	1.0000	Log Likelihood	-1005.13
Standardfehler der Schätzungen (SEy)	297.52	Durbin-Watson (DW) Statistik	1.8588
Anzahl der Beobachtungen	435	Anzahl der Iterationen	5

Die Modelle ARIMA(p,d,q), oder autoregressiver integrierter Prozess des gleitenden Mittelwertes, sind eine Erweiterung des AR-Modells, die drei Komponenten verwenden, um die serielle Korrelation in Zeitreihendaten zu modellieren. Die erste Komponente ist der Begriff Autoregression (AR). Das AR(p) Modell verwendet die p-Verzögerungen der Zeitreihen in der Gleichung. Ein AR(p) Modell trägt die Form: $y(t)=a(1)*y(t-1)+...+a(p)*y(t-p)+e(t)$. Die zweite Komponente ist der Begriff Befehl der Integration (d). Jeder Integrationsbefehl entspricht einer Differenzierung der Zeitreihen. I(1) bedeutet, die Daten einmal zu differenzieren. I(d) bedeutet, die Daten d Male zu differenzieren. Die dritte Komponente ist der Begriff gleitender Mittelwert (MA). Das MA(q) Modell verwendet die q-Verzögerungen von Vorausberechnungsfehlern, um die Vorausberechnung zu verbessern. Ein MA(q) Modell besitzt die Form: $y(t)=e(t)+b(1)*e(t-1)+...+b(q)*e(t-q)$. Zum Schluss, ein ARMA(p,q) Modell besitzt die kombinierte Form: $y(t)=a(1)*y(t-1)+...+a(p)*y(t-p)+e(t)+b(1)*e(t-1)+...+b(q)*e(t-q)$.

Das R-Quadrat, oder Bestimmtheitskoeffizient, zeigt die Prozentvariation in der abhängigen Variablen, die mittels der unabhängigen Variablen in dieser Regressionsanalyse erklärt und begründet werden kann. In einer Mehrfachregression jedoch, berücksichtigt das korrigierte R-Quadrat das Vorhandensein von zusätzlichen unabhängigen Variablen oder Regressoren und korrigiert diesen R-Quadrat-Wert zu einer genaueren Sicht der Regressionserklärungskraft. Allerdings, unter bestimmten Umständen der ARIMA-Modellierung (z.B., mit Modellen ohne Konvergenz), tendiert das R-Quadrat unzuverlässig zu sein.

Der Mehrfachkorrelationskoeffizient (mehrfaches R) misst die Korrelation zwischen der tatsächlichen abhängigen Variablen (Y) und der geschätzten oder angepassten Variablen (Y) basierend auf der Regressionsgleichung. Diese Korrelation ist auch die Quadratwurzel des Bestimmtheitskoeffizienten (R-Quadrat).

Der Standardfehler der Schätzungen (SEy) beschreibt die Dispersion der Datenpunkte über und unter der Regressionslinie oder -ebene. Dieser Wert wird als Teil der Berechnung verwendet, um später das Konfidenzintervall der Schätzungen zu erhalten.

AIC und SC werden oft bei der Modellauswahl verwendet. SC verhängt eine größere Strafe für zusätzliche Koeffizienten. Normalerweise sollte der Benutzer ein Modell mit dem niedrigsten Wert von AIC und SC auswählen.

Die Durbin-Watson-Statistik misst die serielle Korrelation in den Residuen. Normalerweise, eine DW kleiner als 2 impliziert eine positive serielle Korrelation.

Regressionsergebnisse

	Achsabschnitt	AR(1)	MA(1)
Koeffizienten	-0.0626	1.0055	0.4936
Standardfehler	0.3108	0.0006	0.0420
t-Statistik	-0.2013	1691.1373	11.7633
p-Wert	0.8406	0.0000	0.0000
Untere 5%	0.4498	1.0065	0.5628
Obere 95%	-0.5749	1.0046	0.4244

Freiheitsgrade

Freiheitsgrade für die Regression	2
Freiheitsgrade für das Residuum	432
Gesamtfreiheitsgrade	434

Hypothesentest

Kritische t-Statistik (99% Konfidenz mit Freiheitsgrade vor	2.5873
Kritische t-Statistik (95% Konfidenz mit Freiheitsgrade vor	1.9655
Kritische t-Statistik (90% Konfidenz mit Freiheitsgrade vor	1.6484

Die Koeffizienten liefern den geschätzten Achsabschnitt und die geschätzte Steigungen der Regression. Zum Beispiel, die Koeffizienten sind Schätzungen der wahren b-Werten der Bevölkerung in der folgende Regressionsgleichung $Y = b_0 + b_1X_1 + b_2X_2 + ... + b_nX_n$. Der Standardfehler misst wie akkurat die vorausberechneten Koeffizienten sind, und die t-Statistiken sind die Verhältnisse von jedem vorausberechneten Koeffizienten mit seinem Standardfehler.

Die t-Statistik wird beim Hypothesentesten verwendet, wobei man die Nullhypothese (Ho) so einstellt, dass der reelle Mittelwert des Koeffizienten = 0, und die Alternativhypothese (Ha) so einstellt, dass der reelle Mittelwert des Koeffizienten nicht 0 gleicht. Ein t-Test wird ausgeführt und die berechnete t-Statistik wird mit den kritischen Werten an den relevanten Freiheitsgraden für das Residuum verglichen. Der t-Test ist sehr wichtig, da er berechnet, ob jeder der Koeffizienten statistisch signifikant in der Anwesenheit von anderen Regressoren ist. Das bedeutet, dass der t-Test statistisch nachprüft, ob man einen Regressor oder eine unabhängige Variable in der Regression behalten sollte oder nicht.

Der Koeffizient ist statistisch signifikant, wenn seine berechnete t-Statistik die kritische t-Statistik an den relevanten Freiheitsgraden (df) überschreitet. Die drei Hauptkonfidenzniveaus, die verwendet werden, um die Signifikanz zu testen sind 90%, 95% und 99%. Wenn die t-Statistik eines Koeffizienten das kritische Niveau überschreitet, gilt er als statistisch signifikant. Alternativ, der p-Wert berechnet die Ereigniswahrscheinlichkeit jeder t-Statistik, was bedeutet, dass je kleiner der p-Wert, umso mehr signifikant der Koeffizient. Die üblichen Signifikanzniveaus für den p-Wert sind 0,01, 0,05, und 0,10, die den Konfidenzniveaus von 99%, 95% und 99% entsprechen.

Die Koeffizienten mit ihren p-Werten in blau hervorgehoben, zeigen an, dass sie statistisch signifikant beim 90% Konfidenz- oder 0,10 Alphaniveau sind, während die in rot hervorgehobenen anzeigen, dass sie nicht statistisch signifikant bei irgendwelchen anderen Alphaniveaus sind.

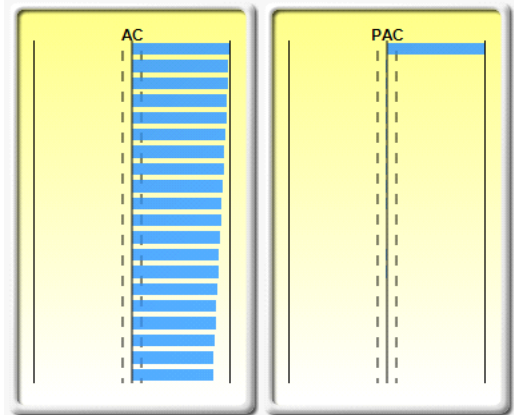
Varianzanalyse

	Summen der Quadrate	Mittelwert der Quadrate	F-Statistik	p-Wert	Hypothesentest
Regression	38415447.53	19207723.76	3171851.1	0.0000	Kritische F-Statistik (99% Konfidenz mit Freiheitsgrade vo
Residuum	2616.05	6.06			Kritische F-Statistik (95% Konfidenz mit Freiheitsgrade vo
Summe	38418063.58				Kritische F-Statistik (90% Konfidenz mit Freiheitsgrade vo

Die Tabelle der Varianzanalyse (ANOVA) bietet einen F-Test der gesamten statistischen Signifikanz des Regressionsmodells. Anstatt die individuellen Regressoren, wie beim t-Test, zu examinieren, betrachtet der F-Test die statistischen Eigenschaften aller geschätzten Koeffizienten. Die F-Statistik wird als das Verhältnis des Mittelwertes der Quadrate der Regression zu dem Mittelwert der Quadrate des Residuums berechnet. Der Zähler misst wie viel von der Regression erklärt wird. Der Nenner indes misst wie viel von der Regression unerklärt bleibt. Deshalb, je größer die F-Statistik, umso mehr signifikant das Modell. Der entsprechende p-Wert wird berechnet, um die Nullhypothese (Ho), wo alle Koeffizienten gleichzeitig Null gleichen, gegen die Alternativhypothese (Ha) zu testen, wo alle Koeffizienten gleichzeitig abweichend von Null sind, hinweisend auf ein signifikantes Gesamtregressionsmodell. Wenn der p-Wert kleiner als die 0,01, 0,05 oder 0,10 Alphasignifikanz ist, dann ist die Regression bedeutsam. Man kann dieselbe Methode bei der F-Statistik anwenden, indem man die berechnete F-Statistik mit den kritischen F Werten an verschiedenen Signifikanzniveaus vergleicht.

Autokorrelation

Zeitverzögerung	AC	PAC	Untere Grenze	Obere Grenze	Q-Statistik	Prob
1	0.9921	0.9921	(0.0958)	0.0958	431.1216	-
2	0.9841	(0.0105)	(0.0958)	0.0958	856.3037	-
3	0.9760	(0.0109)	(0.0958)	0.0958	1,275.4818	-
4	0.9678	(0.0142)	(0.0958)	0.0958	1,688.5499	-
5	0.9594	(0.0098)	(0.0958)	0.0958	2,095.4625	-
6	0.9509	(0.0113)	(0.0958)	0.0958	2,496.1572	-
7	0.9423	(0.0124)	(0.0958)	0.0958	2,890.5594	-
8	0.9336	(0.0147)	(0.0958)	0.0958	3,278.5669	-
9	0.9247	(0.0121)	(0.0958)	0.0958	3,660.1152	-
10	0.9156	(0.0139)	(0.0958)	0.0958	4,035.1192	-
11	0.9066	(0.0049)	(0.0958)	0.0958	4,403.6117	-
12	0.8975	(0.0068)	(0.0958)	0.0958	4,765.6032	-
13	0.8883	(0.0097)	(0.0958)	0.0958	5,121.0697	-
14	0.8791	(0.0087)	(0.0958)	0.0958	5,470.0032	-
15	0.8698	(0.0064)	(0.0958)	0.0958	5,812.4256	-
16	0.8605	(0.0056)	(0.0958)	0.0958	6,148.3694	-
17	0.8512	(0.0062)	(0.0958)	0.0958	6,477.8620	-
18	0.8419	(0.0038)	(0.0958)	0.0958	6,800.9622	-
19	0.8326	(0.0003)	(0.0958)	0.0958	7,117.7709	-
20	0.8235	0.0002	(0.0958)	0.0958	7,428.3952	-



Wenn die Autokorrelation $AC(1)$ ungleich Null ist, bedeutet es, dass die Reihe in der 1. Ordnung seriell korreliert ist. Wenn die $AC(k)$ mehr oder weniger geometrisch mit zunehmender Verzögerung wegstirbt, bedeutet es, dass die Reihe einem autoregressiven Prozess der unteren Ordnung folgt. Wenn die $AC(k)$ nach einer kleinen Anzahl von Verzögerungen auf Null fällt, bedeutet es, dass die Reihe einem gleitenden Mittelwert Prozess der unteren Ordnung folgt. Die partielle Korrelation $PAC(k)$ misst die Korrelation der Werte, die k Perioden auseinander sind, nach Entfernung der Korrelation von den Zwischenverzögerungen. Wenn die Struktur der Autokorrelation von einer Autoregression mit einer Ordnung kleiner als k erfasst werden kann, dann wird die partielle Autokorrelation bei der Verzögerung k nahe Null liegen. Ljung-Box Q-Statistiken und deren p-Werte bei der Verzögerung k weisen die Nullhypothese auf, dass es keine Autokorrelation bis hin zur Ordnung k gibt. Die punktierten Linien in den Diagrammen der Autokorrelationen sind die angenäherten zwei Standardfehlergrenzen. Wenn sich die Autokorrelation innerhalb dieser Grenzen befindet, ist sie nicht signifikant abweichend von Null (zirka) beim 5% Signifikanzniveau.

Vorausberechnung

Periode	tatsächlich (Y)	usberechnung (F)	Fehler (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)

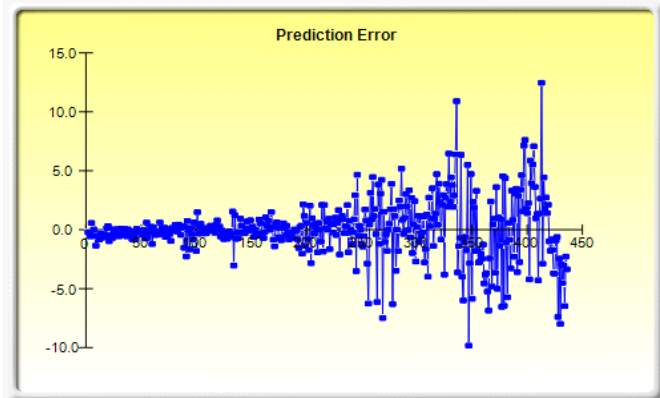
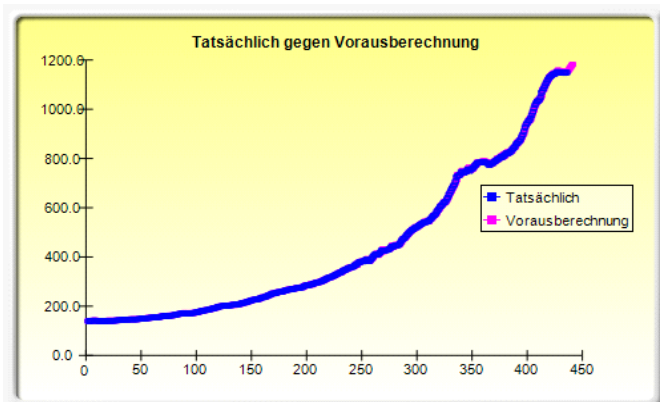


Bild 3.13B - Box Jenkins ARIMA Vorausberechnungsbericht

AUTO ARIMA (Box-Jenkins ARIMA Fortgeschrittene Zeitreihen)

Theorie:

Dieses Tool liefert Analysen, die mit dem ARIMA Modul identisch sind, mit dem Unterschied, dass das Auto-ARIMA Modul Teile der herkömmlichen ARIMA Modellierung automatisiert. Es testet automatisch mehrfache Permutationen von Modellspezifikationen und liefert das bestpassende Modell. Die Ausführung einer Auto-ARIMA ist ähnlich wie bei normalen ARIMA Vorausberechnungen. Der Unterschied ist, dass die P, D, Q Inputs nicht mehr erforderlich sind und dass verschiedene Kombinationen dieser Inputs automatisch ausgeführt und verglichen werden.

Prozedur:

- Excel starten und Ihre Daten eingeben oder ein existierendes vorauszuberechnendes Arbeitsblatt mit historischen Daten öffnen (das im Bild 3.14 angezeigte Beispiel verwendet die Beispielsdatei *Fortgeschrittene Vorausberechnungsmodelle* im Menü *Beispiele* von Risiko Simulator)
- Im Arbeitsblatt *Auto-ARIMA*, wählen Sie *Risiko Simulator | Vorausberechnung | AUTO-ARIMA*. Sie können diese Methode auch folgendermaßen aufrufen: Auf Vorausberechnung im Ikonenband klicken oder irgendwo im Modell rechtsklicken und Vorausberechnung im Kurzbefehlsmenü wählen
- Klicken Sie auf die Ikone Verknüpfung und verknüpfen Sie mit den existierenden Zeitreihendaten. Dann geben Sie die Anzahl der gewünschten Vorausberechnungsperioden ein und klicken Sie auf **OK**

Bemerkung für ARIMA und AUTO-ARIMA:

- Für ARIMA und Auto-ARIMA: Sie können zukünftige Perioden modellieren und berechnen entweder unter Verwendung lediglich der abhängigen Variablen (Y), das heißt, der *Zeitreihenvariable* allein, oder Sie können zusätzliche exogene Variablen ($X_1, X_2, \dots, X_n$) hinzufügen, genau wie in einer Regressionsanalyse wo man mehrfache unabhängige Variablen hat. Wenn Sie nur die Zeitreihenvariable (Y) verwenden, können Sie beliebig viele Vorausberechnungsperioden ausführen. Wenn Sie jedoch exogene Variablen (X) hinzufügen, bemerken Sie bitte, dass die Vorausberechnungsperioden auf die Anzahl der Datenperioden der exogenen Variablen minus den Datenperioden der Zeitreihenvariablen begrenzt sind. Zum Beispiel, Sie können nur bis zu 5 Perioden vorausberechnen, wenn Sie historische Zeitreihendaten von 100 Perioden und nur wenn Sie exogene Variablen von 105 Perioden haben (100 historische Perioden zur Angleichung der Zeitreihenvariablen und 5 zusätzliche Zukunftsperioden von unabhängigen exogenen Variablen, um die abhängige Zeitreihenvariable vorauszuberechnen).

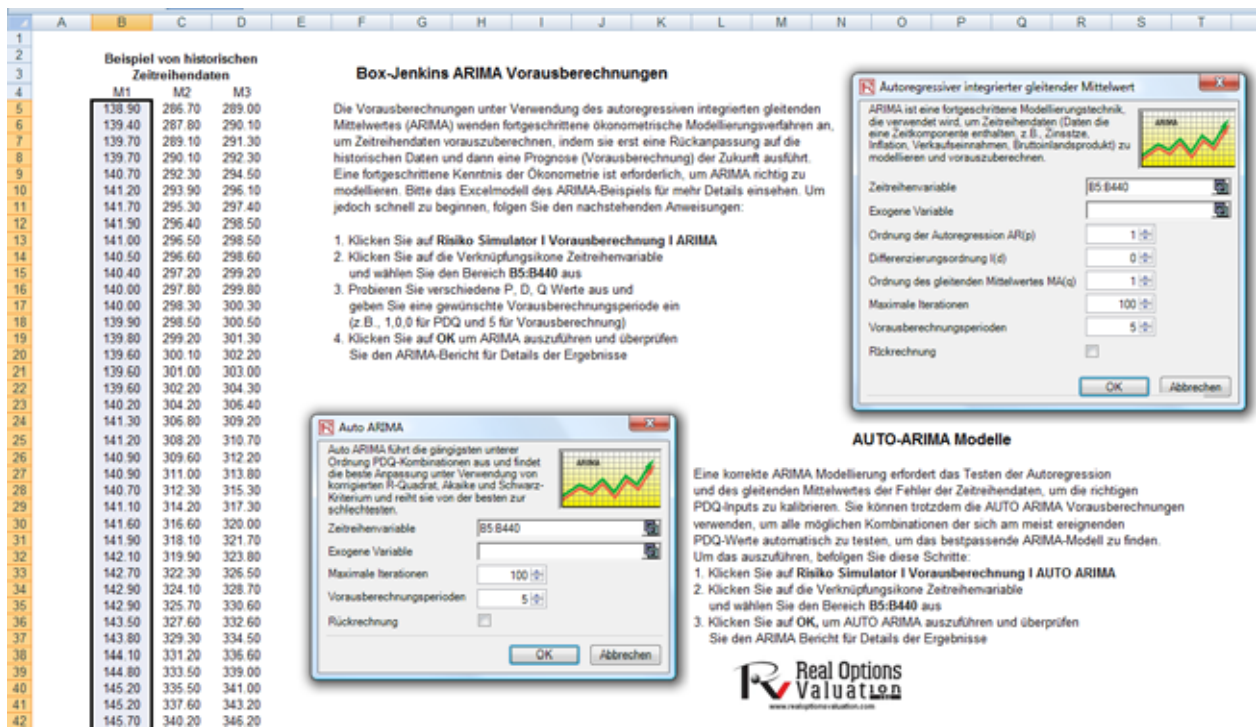


Bild 3.14 - AUTO-ARIMA Modul

Grund-Ökonometrie

Theorie:

Die Ökonometrie bezieht sich auf eine Branche der Geschäftsanalytik: Modellierungs- und Vorausberechnungsverfahren zur Modellierung des Verhaltens oder zur Vorausberechnung von bestimmten geschäftlichen oder wirtschaftlichen Variablen. Die Ausführung der Modelle der Grund-Ökonometrie ist ähnlich wie bei der normalen Regressionsanalyse, außer dass man die abhängigen und unabhängigen Variablen vor der Ausführung einer Regression modifizieren kann. Der generierte Bericht ist der gleiche wie im vorherigen Abschnitt über die mehrfache Regression angezeigt. Auch die Interpretierung ist identisch mit den vorher beschriebenen Interpretierungen.

Prozedur:

- Excel starten und Ihre Daten eingeben oder ein existierendes vorauszuberechnendes Arbeitsblatt mit historischen Daten öffnen (das im Bild 3.15 angezeigte Beispiel verwendet die Beispielsdatei **Fortgeschrittene Vorausberechnungsmodelle** im Menü **Beispiele** von Risiko Simulator)
- Wählen Sie die Daten im Arbeitsblatt **Grund-Ökonometrie** und dann **Risiko Simulator | Vorausberechnung | Grund-Ökonometrie**
- Geben Sie die gewünschten abhängigen und unabhängigen Variablen ein (siehe Bild 3.15 für Beispiele) und klicken Sie auf **OK**, um das Modell und den Bericht auszuführen oder klicken Sie auf **Ergebnisse anzeigen**, um die Ergebnisse vor der Generierung des Berichts anzuschauen, falls Sie Änderungen am Modell durchführen müssen.

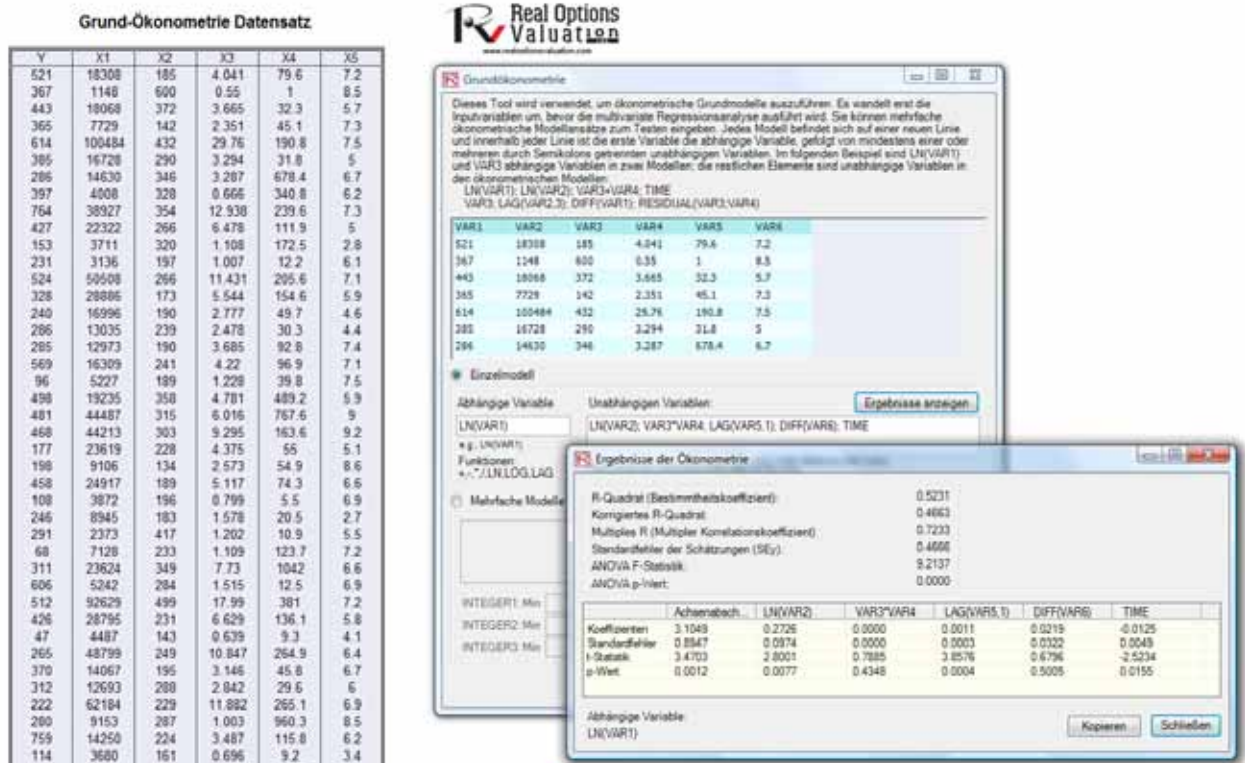


Bild 3.15 - Grund-Ökonometrie Modul

Bemerkung:

- Siehe Kapitel 9 für Details über die Interpretierung von Regressionsoutputs und, im weiteren Sinne, von Outputs einer Grund-Ökonometrie Analyse.
- Um ein ökonometrisches Modell auszuführen, wählen Sie die Daten (B5:G55) einschließlich der Kopfzeilen aus und klicken Sie auf **Risiko Simulator** | **Vorausberechnung** | **Grund-Ökonometrie**. Sie können dann die Variablen und ihre Änderungen für die abhängigen und unabhängigen Variablen eingeben (Bild 8.15). Bitte bemerken Sie, dass nur eine abhängige Variable (Y) erlaubt ist, während mehrfache Variablen im Bereich „unabhängige Variablen“ (X) erlaubt sind, getrennt durch Semikolons ";" und dass man mathematische Grundfunktionen verwenden kann (z.B., LN, LOG, LAG, +, -, /, *, TIME, RESIDUAL, DIFF). Klicken Sie auf **Ergebnisse anzeigen**, um eine Vorschau des berechneten Modells zu bekommen und auf **OK**, um den Bericht des ökonometrischen Modells zu generieren
- Sie können mehrfache Modelle automatisch generieren, indem Sie ein Beispielsmodell eingeben und die vordefinierte „**INTEGER(N)**“ Variable verwenden, ebenso wie durch das wiederholte Verschieben von Daten („**Daten verschieben**“) in spezifische Reihen nach oben oder nach unten. Zum Beispiel, wenn Sie die Variable LAG(VAR1, INTEGER1) verwenden und INTEGER1 zwischen MIN = 1 und MAX = 3 einstellen,

werden die folgenden drei Modelle ausgeführt: LAG(VAR1,1), dann LAG(VAR1,2) und letztlich LAG(VAR1,3) (wobei LAG = Verzögerung). Gelegentlich könnten Sie testen wollen, ob die Zeitreihendaten strukturelle Verschiebungen haben oder ob das Verhalten des Modells konsistent im Laufe der Zeit bleibt, indem Sie die Daten verschieben und dann das gleiche Modell ausführen. Zum Beispiel, wenn Sie 100 chronologisch aufgelistete Monatsdaten haben, können Sie die Daten jeweils drei Monate für 10 Mal nach unten verschieben (das heißt, das Modell wird auf den Monaten 1-100, 4-100, 7-100 und so weiter ausgeführt). Unter Verwendung dieses Bereichs der **mehrfachen Modelle** in der Grund-Ökonometrie, können Sie hunderte von Modellen durch die einfache Eingabe einer einzelnen Modellgleichung ausführen, wenn sie diese Methoden der vordefinierten Ganzzahlvariablen und der Verschiebung verwenden.

J-S Kurven Vorausberechnungen

Theory:

Die J-Kurve oder exponentielle Wachstumskurve ist eine Kurve wo das Wachstum der nächsten Periode vom Niveau der aktuellen Periode abhängt und der Anstieg exponentiell ist. Das heißt, dass im Laufe der Zeit, die Werte von einer Periode zur anderen signifikant wachsen werden. Dieses Modell wird typisch in der Vorausberechnung des biologischen Wachstums und der chemischen Reaktionen im Laufe der Zeit verwendet.

Prozedur:

-  Excel starten und **Risiko Simulator | Vorausberechnung | JS Kurven** wählen.
-  Wählen Sie den J oder S Kurventyp, geben Sie die erforderlichen Inputhypothesen ein (siehe Bilder 3.16 und 3.17 für Beispiele) und klicken Sie auf **OK**, um das Modell und den Bericht auszuführen.

Die S-Kurve oder logistische Wachstumskurve beginnt wie eine J-Kurve mit exponentiellen Wachstumsraten. Im Laufe der Zeit wird die Umgebung gesättigt (z.B., Marktsättigung, Konkurrenz, Überfüllung), das Wachstum lässt nach und der Vorausberechnungswert endet schließlich am Sättigungs- oder Maximalniveau. Dieses Modell wird typischerweise verwendet bei der Vorausberechnung des Marktanteils oder des Verkaufswachstums eines neuen Produktes von der Markteinführung bis zur Reife und den Rückgang, bei der Bevölkerungsdynamik und bei anderen natürlich vorkommenden Phänomenen. Bild 3.17 zeigt ein Beispiel einer S-Kurve.

J-Kurve Exponentialwachstumskurven

In der Mathematik ist eine exponentiell wachsende Menge, eine Menge deren Wachstumsrate immer proportional zu ihrer aktuellen Größe ist. Man sagt, dass ein derartiges Wachstum einem exponentiellen Gesetz folgt. Daraus ergibt sich, dass für eine exponentiell wachsende Menge, je größer die Menge wird, desto schneller wächst sie. Aber es impliziert auch, dass das Verhältnis zwischen der Größe der abhängigen Variablen und ihrer Wachstumsrate durch ein strenges Gesetz der einfachsten Art geregelt wird: die direkte Proportion. Der allgemeine Grundsatz des Exponentialwachstums ist: je größer eine Zahl wird, umso schneller wächst sie. Eine exponentiell wachsende Nummer wird letztendlich größer wachsen als irgendeine andere Nummer, die nur mit einer konstanten Rate über denselben Zeitraum wächst. Diese Vorausberechnungsmethode wird auch die J-Kurve genannt, aufgrund ihrer Form die einem J ähnelt. Diese Wachstumskurve hat kein Maximalniveau. Andere Wachstumskurven sind, unter anderem, die S-Kurven und die Markov-Ketten.

Um eine J-Kurve Vorausberechnung zu generieren, befolgen Sie die nachstehenden Anweisungen:

1. Klicken Sie auf **Risiko Simulator I Vorausberechnung I JS Kurven**
2. Wählen Sie **Exponentielle J-Kurve** aus und geben Sie die gewünschten Inputs ein (z.B., Anfangswert von 100, Wachstumsrate von 5 Prozent, Endperiode von 100)
3. Klicken Sie auf **OK**, um die Vorausberechnung auszuführen und verbringen Sie einige Zeit, den Vorausberechnungsbericht zu überprüfen



Die J-S Kurven stehen für J-Kurve (Exponentialwachstum) und S-Kurve (logistische Wachstumskurve). Diese Kurven werden bei der Vorausberechnung von hohen Wachstumsraten verwendet (J-Kurve) oder auch bei Situationen mit Ereignissen, die ein anfänglich hohes Wachstum aufweisen, das später nachlässt und mit der Zeit reift, wenn die Umgebungskapazität gesättigt ist (S-Kurve).

☒ Exponentielle J-Kurve ☐ Logistische S-Kurve

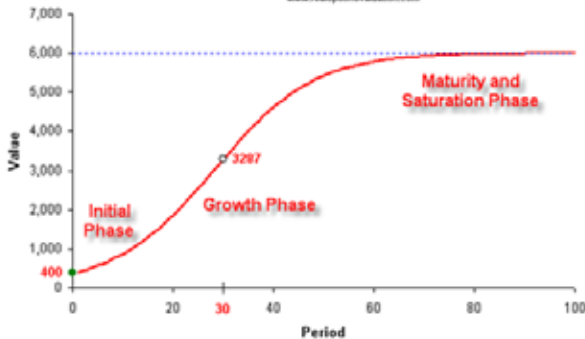
Anfangswert: 100
Wachstumsrate (%): 5
Sättigungsniveau:
Eine Vorausberechnungskurve basierend auf den folgenden Perioden generieren:
Endperiode: 100

OK Abbrechen

Bild 3.16 - J-Kurve Vorausberechnung

Logistische S-Kurve

Eine logistische Funktion oder logistische Kurve modelliert die S-Wachstumskurve irgendeiner Variablen. Das Anfangsstadium des Wachstums ist ungefähr exponentiell, dann, wenn die Konkurrenz auftaucht, verlangsamt sich das Wachstum, um bei der Reifezeit anzuhalten. Diese Funktionen finden Anwendung in zahlreichen Bereichen, von der Biologie bis hin zur Wirtschaftswissenschaft. Zum Beispiel, in der Entwicklung eines Embryos teilt sich ein befruchtetes Ovum und die Zellzahl wächst folgendermaßen: 1, 2, 4, 8, 16, 32, 64, usw. Dies ist exponentielles Wachstum. Aber der Fötus kann nur so groß wie die Größe des Uterus wachsen; daher beginnen andere Faktoren die Steigerung der Zellzahl abzubremsen, und die Wachstumsrate lässt nach (aber das Baby wächst natürlich weiter). Nach einer angemessenen Zeit wird das Kind geboren und wächst weiter. Letztendlich wird die Zellzahl stabil; die Körpergröße der Person ist konstant; das Wachstum wird bei der Reifezeit angehalten. Man kann dieselben Grundsätze auf das Bevölkerungswachstum von Tieren oder Menschen anwenden, und auch bei der Marktdurchdringung und Einnahmen eines Produktes, wo man einen anfänglichen Wachstumsschub in der Marktdurchdringung hat, aber mit der Zeit lässt das Wachstum aufgrund der Konkurrenz nach und letztendlich wird die Marktwachstumsrate sinken und reifen, wenn der Markt gesättigt ist und seine Aufnahmekapazität erreicht hat.



1. Klicken Sie auf **Risiko Simulator I Vorausberechnung I JS Kurven**
2. Geben Sie die erforderlichen Inputs ein (siehe unten für ein Beispiel)
3. Klicken Sie auf **OK** und überprüfen Sie den Vorausberechnungsbericht

Die J-S Kurven stehen für J-Kurve (Exponentialwachstum) und S-Kurve (logistische Wachstumskurve). Diese Kurven werden bei der Vorausberechnung von hohen Wachstumsraten verwendet (J-Kurve) oder auch bei Situationen mit Ereignissen, die ein anfänglich hohes Wachstum aufweisen, das später nachlässt und mit der Zeit reift, wenn die Umgebungskapazität gesättigt ist (S-Kurve).

☐ Exponentielle J-Kurve ☒ Logistische S-Kurve

Anfangswert: 200
Wachstumsrate (%): 10
Sättigungsniveau: 6000
Eine Vorausberechnungskurve basierend auf den folgenden Perioden generieren:
Endperiode: 100

OK Abbrechen

Bild 3.17 - S-Kurve Vorausberechnung

GARCH Volatilitätsvorausberechnungen

Theorie:

Das Modell der verallgemeinerten autoregressiven bedingten Heteroskedastizität (GARCH) wird verwendet, um von einem marktgängigen Wertpapier (z.B., Aktienpreise, Rohstoffpreise, Erdölpreise und so weiter) die historischen Volatilitätsniveaus zu modellieren und die zukünftigen Volatilitätsniveaus vorauszuberechnen. Der Datensatz muss eine Zeitreihe von Rohpreisniveaus sein. Erst konvertiert GARCH die Preise in relative Erträge und dann wird eine interne Optimierung ausgeführt, um die historischen Daten einer zum Mittelwert zurückkehrenden Volatilitätsterminstruktur anzupassen, unter der Annahme, dass die Natur der Volatilität heteroskedastisch ist (sie ändert sich im Laufe der Zeit gemäß einiger ökonometrischer Eigenschaften). Die theoretischen Besonderheiten eines GARCH Modells fallen außerhalb des Rahmens dieses Benutzerhandbuchs. Für mehr Details über GARCH Modelle, lesen Sie bitte *Advanced Analytical Models* von Dr. Johnathan Mun (Wiley 2008).

Prozedur:

- Excel starten und die Beispielsdatei **Fortgeschrittene Vorausberechnungsmodelle** öffnen. Gehen Sie zum Arbeitsblatt **GARCH** und wählen Sie **Risiko Simulator | Vorausberechnung | GARCH**.
- Klicken Sie auf die Ikone Verknüpfung, wählen Sie den *Datenspeicherort* und geben Sie die erforderlichen Inputhypothesen ein (siehe Bild 3.18). Dann klicken Sie auf **OK**, um das Modell und den Bericht auszuführen.

Notiz: Die typische Vorausberechnungssituation der Volatilität erfordert Folgendes: $P = 1$, $Q = 1$, Periodizität = Anzahl der Perioden pro Jahr (12 für Monatsdaten, 52 für Wochendaten, 252 oder 365 für Tagesdaten), Basis = Minimum von 1 und bis zum Periodizitätswert und Vorausberechnungsperioden = Anzahl der annualisierten Volatilitätsvorausberechnungen, die Sie erhalten möchten. Es stehen einige GARCH Modelle in Risiko Simulator zu Verfügung, einschließlich EGARCH, EGARCH-T, GARCH-M, GJR-GARCH, GJR-GARCH-T, IGARCH, und T-GARCH. Siehe *Modeling Risk, Second Edition* (Wiley 2010) über die GARCH Modellierung für mehr Details über die Nutzung jeder Spezifikation.

Historische Daten

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

Um ein GARCH-Modell auszuführen, geben Sie die relevanten Zeitreihendaten ein, dann klicken Sie auf **Risiko Simulator | Vorausberechnung | GARCH**, dann auf die Verknüpfungssikone des Datenspeicherorts und wählen Sie den Bereich der historischen Daten aus (z.B., C8:C2428). Geben Sie die erforderlichen Inputs ein (z.B., P 1, Q 1, tägliche Handelsperiodizität 252, Prädiktive Basis 1, Vorausberechnungsperioden 10) und klicken Sie auf **OK**. Überprüfen Sie den generierten Vorausberechnungsbericht.

Bild 3.18 - GARCH Volatilitätsvorausberechnung

	$z_t \sim \text{Normal}$	$z_t \sim T$
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-Ketten

Theorie:

Eine Markov-Kette existiert, wenn die Wahrscheinlichkeit eines zukünftigen Zustands von einem vorhergehenden Zustand abhängt und wenn zusammengefügt, sie eine Kette bilden, die zurück zu einen Langzeit-Dauerzustandsniveau kehren. Diese Methode wird typisch verwendet, um den Marktanteil von zwei Konkurrenten vorauszuberechnen. Die erforderlichen Inputs sind die Anfangswahrscheinlichkeit, dass ein Kunde des ersten Geschäfts (der erste Zustand) zu diesem selben Geschäft in der nächsten Periode zurückkehren wird, gegen die Wahrscheinlichkeit, dass er in dem nächsten Zustand zum Geschäft eines Konkurrenten wechseln wird.

Prozedur:

-  Excel starten und **Risiko Simulator | Vorausberechnung | Markov-Kette** wählen.
-  Die gewünschten Inputhypothesen eingeben (siehe Bild 3.19 für ein Beispiel) und auf **OK** klicken, um das Modell und den Bericht auszuführen.

Notiz:

Stellen Sie beide Wahrscheinlichkeiten auf 10%, führen Sie die Markov-Kette erneut aus und Sie werden die Effekte des Wechselverhaltens sehr deutlich in das resultierende Bericht sehen.

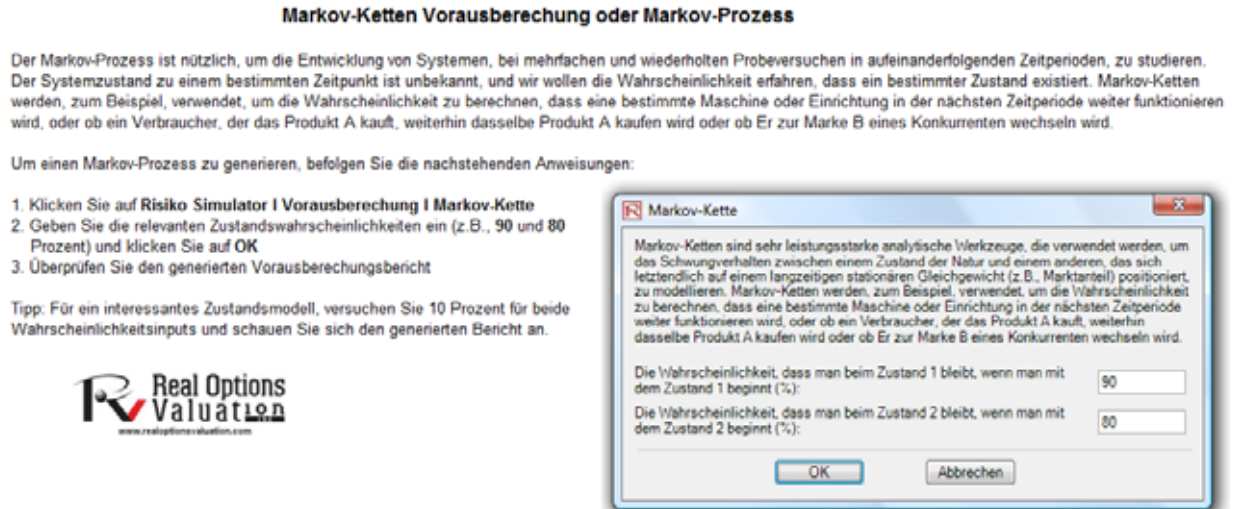


Bild 3.19 - Markov-Ketten (Wechselregime)

Maximale Wahrscheinlichkeitsmodelle (MLE) auf Logit, Probit und Tobit

Theorie:

„Begrenzte abhängige Variablen“ beschreibt die Situation in der die abhängige Variable, Daten mit begrenztem Geltungsbereich und Umfang, wie binäre Antworten (0 oder 1), gestutzte, geordnete oder zensierte Daten, enthält. Zum Beispiel, gegeben ein Satz von unabhängigen Variablen (z.B., Alter, Einkommen, Bildungsgrad von Kreditkarteninhabern oder Hypothekennehmern), kann man die Ausfallwahrscheinlichkeit unter Verwendung der Schätzung der maximalen Wahrscheinlichkeit (MLE) modellieren. Die Antwort oder abhängige Variable Y ist binär, das heißt, dass sie nur zwei mögliche Ergebnisse haben kann, die wir als 1 und 0 bezeichnen (z.B., Y könnte die Anwesenheit/Abwesenheit eines bestimmten Zustands repräsentieren, Rückzahlung/nicht Rückzahlung von vergangenen Darlehen, Erfolg/Misserfolg eines bestimmten Gerätes, Ja/Nein Antwort einer Umfrage, usw.). Außerdem haben wir einen Vektor von unabhängigen Variablenregressoren X, von denen man annimmt, dass sie das Ergebnis Y beeinflussen. Eine typisch gewöhnliche Regressionsmethode der kleinsten Quadrate ist ungültig, weil die Regressionsfehler heteroscedastisch und nicht-normal sind und die resultierenden geschätzten Wahrscheinlichkeitsschätzungen werden unsinnige Werte über 1 oder unter 0 ergeben. Die MLE-Analyse behandelt diese Probleme unter Verwendung einer iterativen Optimierungsroutine, um eine logarithmische Wahrscheinlichkeitsfunktion zu maximieren, wenn die abhängigen Variablen begrenzt sind.

Logit oder logistische Regression wird verwendet, um die Auftrittswahrscheinlichkeit eines Ereignisses, durch die Anpassung von Daten zu einer logistischen Kurve, vorzuberechnen. Es ist ein verallgemeinertes für Binomregressionen verwendetes Linearmodell und sowie viele Formen der Regressionsanalyse, verwendet es einige Prädiktorvariablen, die entweder numerisch oder kategorisch sein könnten. MLE angewendet in einer binären multivariaten logistische Analyse wird verwendet, um bestimmte abhängige Variablen zu modellieren, um die erwartete Erfolgswahrscheinlichkeit der Zugehörigkeit zu einer bestimmten Gruppe zu bestimmen. Die geschätzten Koeffizienten für das Logitmodell sind die logarithmischen Chancenverhältnisse und können nicht direkt als Wahrscheinlichkeiten interpretiert werden. Erst ist eine schnelle Berechnung erforderlich und die Methode ist einfach.

Im Besonderen, das Logitmodell ist bezeichnet als geschätztes $Y = \text{LN}[P_i/(1-P_i)]$ oder umgekehrt $P_i = \text{EXP}(\text{geschätztes } Y)/(1+\text{EXP}(\text{geschätztes } Y))$ und die Koeffizienten β_i sind die logarithmischen Chancenverhältnisse, sodass wenn man den Antilogarithmus oder $\text{EXP}(\beta_i)$ nimmt, erhält man die Chancenverhältnisse von $P_i/(1-P_i)$. Das bedeutet, dass mit einer Steigerung in einer Einheit von β_i , wächst das Chancenverhältnis um diese Menge. Letztlich, die Änderungsrate in der Wahrscheinlichkeit $dP/dX = \beta_i P_i(1-P_i)$. Der Standardfehler misst wie akkurat die prognostizierten Koeffizienten sind. Die t-Statistiken sind die Verhältnisse jedes vorausgerechneten Koeffizienten zu seinem Standardfehler und werden im typischen Regressionshypothesentest der Signifikanz jedes geschätzten Parameters verwendet. Um die Erfolgswahrscheinlichkeit der Zugehörigkeit zu einer bestimmten Gruppe zu schätzen (z.B., die Prognostizierung ob ein Raucher Lungenkomplikationen entwickeln wird, gegeben einer gerauchten Menge pro Jahr), einfach den geschätzten Y-Wert berechnen, unter Verwendung der Koeffizienten der Schätzung der größten Wahrscheinlichkeit (MLE). Zum Beispiel, mit einem Modell, wobei $Y = 1.1 + 0.005$ (Zigaretten), dann wird jemand der/die 100 Packungen pro Jahr raucht einen geschätzten Y-Wert von $1.1 + 0.005(100) = 1.6$ haben. Dann berechnet man den inversen Antilogarithmus des Chancenverhältnisses wie folgt: $\text{EXP}(\text{geschätztes } Y)/[1 + \text{EXP}(\text{geschätztes } Y)] = \text{EXP}(1.6)/(1+\text{EXP}(1.6)) = 0.8320$. Daraus erfolgt, dass diese Person eine Chance von 83.20% hat, Lungenkomplikationen während seines Lebens zu entwickeln.

Ein Probitmodell (gelegentlich auch als Normitmodell bekannt) ist eine gängige alternative Spezifikation für ein Binärantwortmodell, das eine Probitfunktion verwendet, die unter Verwendung der Schätzung der maximalen Wahrscheinlichkeit geschätzt wird. Die Methode wird Probitregression genannt. Die Probit- und logistischen Regressionsmodelle neigen zur Produktion von sehr ähnlichen Voraussagen, wobei die Parameterschätzungen in einer logistischen Regression dazu tendieren 1.6 bis 1.8 Male höher als in einem entsprechenden Probitmodell zu sein. Die Entscheidung der Verwendung eines Probit- oder Logitmodells hängt gänzlich von der Zweckmäßigkeit ab und der Hauptunterschied liegt darin, dass die logistische Verteilung eine höhere Kurtosis (dickere Schwänze) hat, um Extremwerte zu erklären. Zum Beispiel, angenommen dass Wohnungseigentum die zu modellierende Entscheidung ist und dass diese Antwortvariable binär (Wohnungskauf oder kein Wohnungskauf) ist und von einer Serie von

unabhängigen Variablen X_i , wie Einkommen, Alter, und so weiter, abhängt, sodass $I_i = \beta_0 + \beta_1 X_1 + \dots + \beta_n X_n$, wobei je größer der Wert von I_i , desto höher die Wahrscheinlichkeit des Wohnungseigentums. Für jede Familie gibt es eine kritische Schwelle I^* . Wenn diese überschritten wird, wird die Wohnung gekauft, andernfalls wird keine Wohnung gekauft. Es wird angenommen, dass die Ergebniswahrscheinlichkeit (P) normal verteilt ist, sodass $P_i = \text{CDF}(I)$ unter Verwendung einer standardnormalen kumulativen Verteilungsfunktion (CDF) ist. Daher verwenden Sie die geschätzten Koeffizienten genau so wie die eines Regressionsmodells und unter Verwendung des geschätzten Y -Wertes, wenden Sie eine Standardnormalverteilung an (Sie können entweder die Excelfunktion NORMSDIST oder das Analysetool des Risiko Simulators verwenden, indem Sie die Normalverteilung auswählen und den Mittelwert auf 0 und die Standardabweichung auf 1 einstellen). Zum Schluss, um Probit oder das Wahrscheinlichkeitseinheitsmaß zu erhalten, stellen Sie $I_i + 5$ ein (das liegt daran, dass wann immer die Wahrscheinlichkeit $P_i < 0.5$ ist, ist das geschätzte I_i negativ, auf Grund der Tatsache, dass die Normalverteilung symmetrisch um einen Mittelwert von Null ist).

Das Tobitmodell (zensiertes Tobit) ist eine ökonometrische und biometrische Modellierungsmethode, die verwendet wird, um das Verhältnis zwischen einer nicht negativen abhängigen Variablen Y_i und einer oder mehreren unabhängigen Variablen X_i zu beschreiben. Ein Tobitmodell ist ein ökonometrisches Modell in dem die abhängige Variable zensiert ist; das heißt, die abhängige Variable ist zensiert, weil Werte unter Null nicht beobachtet werden. Das Tobitmodell nimmt an, dass es eine latente nicht beobachtbare Variable Y^* gibt. Diese Variable ist auf den X_i Variablen linear abhängig über einen Vektor von β_i Koeffizienten, die ihre Zwischenverhältnisse bestimmen. Es gibt zusätzlich einen normalverteilten Fehlerbegriff U_i , um Zufallseinflüsse auf dieses Verhältnis zu erfassen. Die beobachtbare Variable Y_i wird als gleich der latenten Variablen definiert, wenn immer die latenten Variablen über Null sind, sonst wird angenommen, dass Y_i Null ist. Das heißt, $Y_i = Y^*$ wenn $Y^* > 0$ und $Y_i = 0$ wenn $Y^* \leq 0$. Wenn der Verhältnisparameter β_i unter Verwendung der normalen kleinsten Quadratenregression des beobachteten Y_i auf X_i geschätzt wird, sind die resultierenden Regressionsschätzer inkonsistent und ergeben nach unten verzerrten Steigungskoeffizienten und einen nach oben verzerrten Achsabschnitt. Nur MLE würde für einen Tobitmodell konsistent sein. In dem Tobitmodell gibt es eine untergeordnete Statistik genannt Sigma, die dem Standardfehler der Schätzung in einer standardnormalen kleinsten Quadratenregression äquivalent ist und die geschätzten Koeffizienten werden genau so wie eine Regressionsanalyse verwendet.

Prozedur:

-  Excel starten und die Beispielsdatei **Fortgeschrittene Vorausberechnungsmodelle** öffnen. Gehen Sie zum Arbeitsblatt **MLE**, wählen Sie den Datensatz einschließlich der Kopfzeilen aus und klicken Sie auf **Risiko Simulator | Vorausberechnung | Maximale Wahrscheinlichkeit**.
-  Wählen sie die abhängige Variable aus der Dropdownliste (siehe Bild 3.20) und klicken Sie auf **OK**, um das Modell und den Bericht auszuführen.

Binäre logistische maximale Wahrscheinlichkeits-Vorausberechnung: LOGIT, PROBIT, TOBIT

LOGIT & PROBIT

Notleidende	Alter	Bildungs- grad	Jahre mit dem aktuellen Arbeitgeber	Jahre an der aktuellen Adresse	Haushalts- einkommen (tausende \$)	Verschuldung/ Einkommen Verhältnis (%)	Kreditkarten- verschuldung (tausende \$)	Sonstige Verschuldung (tausende \$)
1	41	3	17	12	176	9.3	11.36	5.01
0	27	1	10	6	31	17.3	1.36	4
0	40	1	15	14	55	5.5	0.86	2.17
0	41	1	15	14	120	2.9	2.66	0.82
1	24	2	2	0	28	17.3	1.79	3.06
0	41	2	5	5	25	10.2	0.39	2.16
0	39	1	20	9	67	30.6	3.83	16.67
0	43	1	12	11	38	3.6	0.13	1.24
1	24	1	3	4	19	24.4	1.36	3.28
0	36	1						
0	27	1						
0	25	1						
0	52	1						
0	37	1						
0	48	1						
1	36	2						
1	36	2						
0	43	1						
0	39	1						
0	41	3						
0	39	1						
0	47	1						
0	28	1						
0	29	1						
1	21	2						
0	25	4						
0	45	2						
0	43	1						
0	33	2						
0	26	3						
0	45	1						
0	30	1						
0	27	3						
0	25	1						
0	25	1						
			8	4	27	14.4	1.02	2.87
			8	1	35	2.9	0.08	0.94

Bild 3.20 - Maximale Wahrscheinlichkeit Modul

Spline (Kubischer Spline Interpolation und Extrapolation)

Theorie:

Gelegentlich gibt es fehlende Werte in einem Zeitreihendatensatz. Man hat, zum Beispiel, die Zinssätze für die Jahre 1 bis 3, gefolgt von den Jahren 5 bis 8 und dann für Jahr 10. Man kann Spline-Kurven verwenden, um die Zinssatzwerte der fehlenden Jahre basierend auf den existierenden Daten zu interpolieren. Spline-Kurven können auch verwendet werden, um Werte von zukünftigen Perioden über die Zeitperiode der existierenden Daten hinaus vorzuberechnen oder zu extrapolieren. Die Daten können linear oder nichtlinear sein. Bild 3.21 zeigt wie man einen kubischen Spline ausführt und Bild 3.22 zeigt den resultierenden Vorausberechnungsbericht dieses Moduls. Die bekannte X Werte repräsentieren die Werte auf der x-Achse eines Diagramms (in unserem Beispiel sind es die Jahre der bekannten Zinssätze;

normalerweise enthält die x- Achse, die im Voraus bekannten Werte, wie Zeit oder Jahre) und die bekannten Y Werte repräsentieren die Werte auf der y-Achse (in unserem Fall sind es die bekannten Zinssätze). Die Variable der y-Achse ist typischerweise die Variable, von der Sie fehlenden Werte interpolieren oder zukünftige Werte extrapolieren möchten.

Kubische Spline Interpolation und Extrapolation

Das Modell der kubischen Spline Polynominterpolation und -extrapolation wird verwendet, um die „Lücken“ von fehlenden Spotrenditen und Zinssatzterminstrukturen „auszufüllen“, wobei man das Modell verwenden kann, sowohl um fehlende Datenpunkte innerhalb einer Zeitreihe von Zinssätzen (so wie auch makroökonomische Variablen wie Inflationsraten und Rohstoffpreise oder Markteinkünfte) zu interpolieren als auch außerhalb des gegebenen oder bekannten Bereiches zu extrapolieren, was es nützlich für Vorausberechnungszwecke macht.

Jahre	Spotrenditen
0.0833	4.55%
0.2500	4.47%
0.5000	4.52%
1.0000	4.39%
2.0000	4.13%
3.0000	4.16%
5.0000	4.26%
7.0000	4.38%
10.0000	4.56%
20.0000	4.88%
30.0000	4.84%

Dies sind bekannte Renditen und werden als Inputs im Modell der kubischen Spline Interpolation und Extrapolation verwendet.

Um die kubische Spline Vorausberechnung auszuführen, klicken Sie auf **Risiko Simulator | Vorausberechnung | Kubische Spline** und dann auf die Verknüpfungssikone. Wählen Sie C15:C25 als die bekannten X-Werte (Werte auf der x-Achse eines Zeitreihendiagramms) und D15:D25 als die bekannten Y-Werte (achten Sie darauf, dass die Länge der bekannten X- und Y-Werte dieselbe ist). Geben Sie die gewünschten Vorausberechnungsperioden ein (z.B., Anfang 1, Ende 50, Schrittgröße 0.5). Klicken Sie auf OK und prüfen Sie die generierten Vorausberechnungen und das generierte Diagramm nach.

Kubische Spline

Das Modell der kubischen Spline Polynominterpolation und -extrapolation wird verwendet, um die „Lücken“ von fehlenden Werten „auszufüllen“ und um Zeitreihendaten vorzuberechnen. Demzufolge kann man das Modell verwenden, sowohl um fehlende Datenpunkte innerhalb einer Datenzeitreihe (z.B., Ertragskurven, Zinssätze, Makroökonomische Variablen wie Inflationsraten und Rohstoffpreise oder Markteinkünfte) zu interpolieren als auch außerhalb des gegebenen oder bekannten Bereiches zu extrapolieren, was es nützlich für Vorausberechnungen macht.

Bekannte X-Werte: C15:C25

Bekannte Y-Werte: D15:D25

Eine Splinekurve basierend auf den folgenden X-Werten kreieren.

Anfang: 1 Ende: 50 Schrittgröße: 0.5

OK Abbrechen

Bild 3.21 - Kubischer Spline Modul

Prozedur:

- Excel starten und die Beispielsdatei **Fortgeschrittene Vorausberechnungsmodelle** öffnen. Gehen Sie zum Arbeitsblatt **Kubischer Spline**, wählen Sie den Datensatz einschließlich der Kopfzeilen aus und klicken Sie auf **Risiko Simulator | Vorausberechnung | Kubischer Spline**.
- Wenn Sie die Daten zuerst auswählen, wird der Datenspeicherort automatisch in die Bedienungsoberfläche eingefügt. Sie können aber auch manuell auf die Ikone Verknüpfung klicken und die *bekannten X Werte* und *bekannten Y Werte* verknüpfen (siehe Bild 3.21 für ein Beispiel). Dann geben Sie die zur Extrapolierung und Interpolation erforderlichen Werte für *Anfang* und *Ende*, sowie auch die erforderliche *Schrittgröße* zwischen diesen Anfangs- und Endwerten ein. Klicken Sie auf **OK**, um das Modell und den Bericht auszuführen (siehe Bild 3.22).

Kubische Spline-Voraberechnungen

Das Modell der kubischen Spline Polynominterpolation und -extrapolation wird verwendet, um die „Lücken“ von fehlenden Werten „auszufüllen“ und um Zeitreihendaten vorzuberechnen. Demzufolge kann man das Modell verwenden, sowohl um fehlende Datenpunkte innerhalb einer Datenzeitreihe (z.B., Ertragskurven, Zinssätze, Makroökonomische Variablen wie Inflationsraten und Rohstoffpreise oder Markteinkünfte) zu interpolieren als auch außerhalb des gegebenen oder bekannten Bereiches

Ergebnisse der Spline Interpolation und Extrapolation

X	Angepasstes Y	Notizen
1.0	4.39%	Interpolieren
1.5	4.21%	Interpolieren
2.0	4.13%	Interpolieren
2.5	4.13%	Interpolieren
3.0	4.16%	Interpolieren
3.5	4.19%	Interpolieren
4.0	4.22%	Interpolieren
4.5	4.24%	Interpolieren
5.0	4.26%	Interpolieren
5.5	4.29%	Interpolieren
6.0	4.32%	Interpolieren
6.5	4.35%	Interpolieren
7.0	4.38%	Interpolieren
7.5	4.41%	Interpolieren
8.0	4.44%	Interpolieren
8.5	4.47%	Interpolieren
9.0	4.50%	Interpolieren
9.5	4.53%	Interpolieren
10.0	4.56%	Interpolieren
10.5	4.59%	Interpolieren
11.0	4.61%	Interpolieren
11.5	4.64%	Interpolieren
12.0	4.66%	Interpolieren

Dies sind die bekannten Wertinputs in das Modell der kubischen Spline Interpolation und Extrapolation:

Beobachtung	Bekanntes X	Bekanntes Y
1	0.0833	4.55%
2	0.2500	4.47%
3	0.5000	4.52%
4	1.0000	4.39%
5	2.0000	4.13%
6	3.0000	4.16%
7	5.0000	4.26%
8	7.0000	4.38%
9	10.0000	4.56%
10	20.0000	4.88%
11	30.0000	4.84%

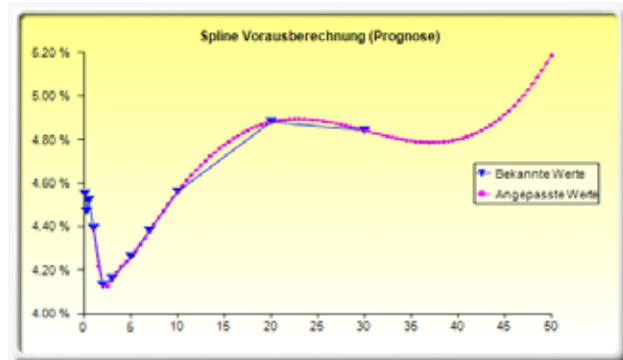


Bild 3.22 - Ergebnisse der Spline-Voraberechnung

4. OPTIMIERUNG

Dieser Abschnitt examiniert die Optimierungsprozesse und -methodologien in ausführlicherem Detail bezüglich ihrer Verwendung im Risiko Simulator. Diese Methodologien schließen die Verwendung der kontinuierlichen gegen die diskreten ganzzahligen Optimierungen, sowie auch der statischen gegen die dynamischen und stochastischen Optimierungen ein.

Optimierungsmethodologien

Es existieren viele Algorithmen, um die Optimierung auszuführen und es sind viele verschiedene Verfahren vorhanden, wenn die Optimierung mit der Monte-Carlo-Simulation verbunden wird. In Risiko Simulator gibt es drei verschiedenartige Optimierungsverfahren and Optimierungstypen, sowie auch unterschiedliche Typen von Entscheidungsvariablen. Risiko Simulator kann, zum Beispiel, Folgendes bearbeiten: **Kontinuierliche Entscheidungsvariablen** (1.2535, 0.2215 und so weiter), **Ganzzahlige Entscheidungsvariablen** (z.B., 1, 2, 3, 4 oder 1.5, 2.5, 3.5 und so weiter), **Binäre Entscheidungsvariablen** (1 und 0 für Go/No-Go Entscheidungen) und **Gemischte Entscheidungsvariablen** (sowohl ganzzahlige und kontinuierliche Variablen). Risiko Simulator kann außerdem **Lineare Optimierungen** (das heißt, wenn sowohl das Ziel als auch die Einschränkungen alle lineare Gleichungen und Funktionen sind) und **Nicht lineare Optimierungen** (das heißt, wenn das Ziel und die Einschränkungen eine Mischung von linearen und nicht linearen Funktionen und Gleichungen sind) bearbeiten.

Bezüglich des Optimierungsverfahrens, kann man Risiko Simulator verwenden, um eine **Diskrete Optimierung** auszuführen. Diese ist eine Optimierung, die auf einem diskreten oder statischen Modell ausgeführt wird, wobei keine Simulationen durchgeführt werden. In anderen Worten, alle Inputs des Modells sind statisch und gleichbleibend. Dieser Optimierungstyp ist anwendbar, wenn das Modell als bekannt angenommen wird und wenn keine Ungewissheiten vorhanden sind. Man kann auch eine diskrete Optimierung zuerst ausführen, um ein optimales Portfolio und seine entsprechende optimale Aufteilung der Entscheidungsvariablen zu bestimmen, bevor man fortgeschrittenere Optimierungsprozeduren anwendet. Zum Beispiel, bevor man eine stochastische Optimierungsaufgabe ausführt, wird eine diskrete Optimierung zuerst durchgeführt, um zu bestimmen, ob es Lösungen zur Optimierungsaufgabe gibt, bevor man eine langwierigere Analyse ausführt.

Die **Dynamische Optimierung** wird angewendet, wenn eine Monte-Carlo-Simulation zusammen mit einer anderen Optimierung ausgeführt wird. Ein anderer Name für solch eine Prozedur ist **Simulations-Optimierung**. Das heißt, zuerst wird eine Simulation ausgeführt, dann werden die Ergebnisse der Simulation im Excelmodell angewendet und dann wird eine Optimierung auf die simulierte Werten angewendet. In anderen Worten, zuerst wird eine Simulation mit N Probeversuchen ausgeführt und dann wird ein Optimierungsverfahren für M Iterationen durchgeführt, bis man die optimalen Ergebnisse oder einen unausführbaren Satz findet. Das heißt,

unter Verwendung des Optimierungsmoduls von Risiko Simulators können Sie auswählen, welche Vorausberechnungs- und Hypothesenstatistiken Sie nach dem Simulationslauf verwenden und im Modell ersetzen möchten. Dann können diese Vorausberechnungsstatistiken in dem Optimierungsverfahren angewendet werden. Diese Methode ist nützlich, wenn Sie ein großes Modell mit vielen aufeinander wirkenden Hypothesen und Vorausberechnungen haben und wenn einige der Vorausberechnungsstatistiken in der Optimierung erforderlich sind. Wenn, zum Beispiel, die Standardabweichung einer Hypothese oder Vorausberechnung im Optimierungsmodell erforderlich ist (z.B., Berechnen des Sharpe-Ratio bei Assetallokation und Optimierungsaufgaben, wobei wir den Mittelwert geteilt durch die Standardabweichung des Portfolios haben), dann sollte man diese Methode anwenden.

Die **Stochastische Optimierung**, dagegen, ähnelt dem dynamischen Optimierungsverfahren, mit der Ausnahme, dass das gesamte dynamische Optimierungsverfahren für T Male wiederholt wird. Das heißt, eine Simulation mit N Probeversuchen wird ausgeführt und dann wird ein Optimierungsverfahren für M Iterationen durchgeführt, um die optimalen Ergebnisse zu erhalten. Dann wird der Prozess T Male repliziert. Die Ergebnisse sind ein Vorausberechnungsdiagramm jeder Entscheidungsvariablen mit T Werten. In anderen Worten, es wird eine Simulation ausgeführt und die Vorausberechnungs- oder Hypothesenstatistiken werden im Optimierungsmodell verwendet, um die optimale Aufteilung der Entscheidungsvariablen zu finden. Dann wird eine weitere Simulation ausgeführt, die andere Vorausberechnungsstatistiken generiert und diese neuen aktualisierten Werte werden dann optimiert, und so weiter. Deshalb hat jede Entscheidungsvariable ihr eigenes Vorausberechnungsdiagramm, das den Bereich der optimalen Entscheidungsvariablen anzeigt. Zum Beispiel, anstatt Einzelpunktschätzungen in dem dynamischen Optimierungsverfahren zu erhalten, können Sie jetzt eine Verteilung der Entscheidungsvariablen bekommen und, daher, einen Bereich von optimalen Werten für jede Entscheidungsvariable. Dies ist auch als eine stochastische Optimierung bekannt.

Zum Schluss, ein „Effiziente Grenze“ Optimierungsverfahren wendet die Begriffe der marginalen Inkremente und der Schattenpreisfindung in Optimierungen an. Das heißt, was würde mit den Ergebnissen der Optimierung geschehen, wenn eine der Einschränkungen leicht gelockert wäre? Sagen wir mal zum Beispiel, wenn die Budgeteinschränkung bei \$1 Millionen festgelegt wäre. Was würde mit dem Ausgang und den optimalen Entscheidungen des Portfolios geschehen, wenn die Einschränkungen jetzt auf \$1.5 Millionen oder \$2 Millionen, und so weiter, festgelegt wären. Dies ist der Begriff der effizienten Grenzen von Markowitz in der Investitionsfinanz: Wenn man der Standardabweichung des Portfolios erlaubt leicht zu steigen, welche zusätzlichen Erträge würde das Portfolio generieren? Dieser Prozess ähnelt dem dynamischen Optimierungsverfahren, mit der Ausnahme, dass eine der Einschränkungen sich ändern darf und dass die Simulation und das Optimierungsverfahren bei jeder Änderung durchgeführt wird. Dieses Verfahren wird am besten manuell unter Verwendung von Risiko Simulator angewendet. Das heißt, zuerst eine dynamische oder stochastische Optimierung durchführen, dann eine weitere Optimierung mit einer Einschränkung erneut ausführen und schließlich diese Prozedur mehrere Male wiederholen.

Dieser manuelle Prozess ist wichtig, weil durch die Änderung der Einschränkung, der Analyst feststellen kann, ob die Ergebnisse ähnlich oder abweichend sind und, deshalb, ob sich eine zusätzliche Analyse lohnt, oder auch feststellen, wie groß der marginale Anstieg in der Einschränkung sein sollte, um eine signifikante Änderung im Ziel und in den Entscheidungsvariablen zu erhalten.

Ein Punkt ist beachtenswert. Es existieren andere Softwareprodukte, die angeblich stochastische Optimierungen durchführen, aber in Wirklichkeit tun sie das nicht. Zum Beispiel, Folgendes ist eine Zeit- und Ressourcenverschwendung: Nach der Ausführung einer Simulation, wird *eine erste* Iteration des Optimierungsverfahren generiert, dann wird wieder eine Simulation ausgeführt und eine *zweite* Optimierungsiteration generiert und so weiter. In anderen Worten, in einer Optimierung wird das Modell durch einen strengen Satz von Algorithmen geführt, wobei mehrfache Iterationen (von einigen bis zu Tausenden von Iterationen) erforderlich sind, um die optimalen Ergebnisse zu erhalten. Daher ist die Generierung von *einer* Iteration nach der anderen eine Zeit- und Ressourcenverschwendung. Man kann das gleiche Portfolio unter Verwendung des Risiko Simulators in einer Minute lösen im Vergleich zu mehreren Stunden bei Verwendung solch einer rückständigen Methode. Außerdem liefert solch eine Simulations-Optimierungsmethode typischerweise schlechte Ergebnisse und ist nicht eine stochastische Optimierungsmethode. Wenn Sie eine Optimierung auf Ihrem Modell ausführen, hüten Sie sich sehr vor diesen Methodologien.

Es folgen zwei Beispiele von Optimierungsaufgaben. Eine verwendet die kontinuierlichen Entscheidungsvariablen, während die andere diskrete ganzzahlige Entscheidungsvariablen verwendet. In beiden Modellen kann man eine diskrete Optimierung, eine dynamische Optimierung, eine stochastische Optimierung oder sogar die effizienten Grenzen mit Schattenpreisfindung anwenden. Man kann irgendeine dieser Methoden für diese beiden Beispiele verwenden. Der Einfachheit halber wird deshalb nur die Modellaufstellung erläutert und die Entscheidung über welches Optimierungsverfahren ausgeführt werden soll wird dem Benutzer überlassen. Ferner verwendet das kontinuierliche Modell die nicht lineare Optimierungsmethode (das liegt daran, dass das berechnete Portfoliorisiko eine nicht lineare Funktion ist und das Ziel eine nicht lineare Funktion der Portfolioerträge geteilt durch die Portfoliorisiken ist), während das zweite Beispiel einer Ganzzahloptimierung, ein Beispiel eines linearen Optimierungsmodells ist (ihr Ziel und alle Einschränkungen sind linear). Infolgedessen enthalten diese beiden Beispiele alles Wesentliche der obengenannten Verfahren.

Optimierung mit kontinuierlichen Entscheidungsvariablen

Bild 4.1 erläutert das Beispiel eines kontinuierlichen Optimierungsmodells. Dieses Beispiel verwendet die Datei *Kontinuierliche Optimierung*, die entweder im Startmenü unter **Start | Real Options Valuation | Risiko Simulator | Beispiel** oder direkt unter **Risiko Simulator | Beispielsmodelle** zu finden ist. In dieses Beispiel gibt es 10 individuelle Aktivaklassen (z.B., verschiedene Typen von Anlagefonds, Aktien oder Aktiva), wobei die Idee ist, den Portfoliobestand meisteffizient und -wirksam aufzuteilen, sodass man den besten „Bang-for-the-Buck“ (Gewinn) erhält. Das heißt, um die bestmöglichen Portfoliorenditen zu generieren, gegeben den in jeder Aktivaklasse enthaltenen Risiken. Um den Begriff der Optimierung wirklich zu begreifen, müssen wir dieses Beispielsmodell ausführlicher erforschen, um zu sehen wie man das Optimierungsverfahren am besten anwenden kann.

Das Modell zeigt die 10 Aktivaklassen und jede Aktivaklasse hat seinen eigenen Satz von annualisierten Renditen und Volatilitäten. Diese Renditen- und Risikomessungen sind annualisierte Werte, sodass man sie konsequent über verschiedene Aktivaklassen vergleichen kann. Renditen werden unter Verwendung des geometrischen Mittels der relativen Renditen berechnet, während die Risiken unter Verwendung der Methode der logarithmischen relativen Aktienrenditen berechnet werden. Siehe den Anhang in diesem Kapitel für Details über die Berechnung der annualisierten Volatilität und der annualisierten Renditen einer Aktie oder Aktivaklasse.

	A	B	C	D	E	F	G	H	I	J	K	L
	1											
	2											
	3											
	4											
		OPTIMIERUNGSMODELL FÜR DIE AUFTEILUNG DER AKTIVA										
		Beschreibung der Aktivaklasse	Annualisierte Renditen	Volatilitäts- risiko	Aufteilungs- gewichte	Erforderliche minimale Aufteilung	Erforderliche maximale Aufteilung	Renditen- Risiko- Verhältnis	Renditen- Rang (hoch- niedrig)	Risiko- Rang (hoch- niedrig)	Renditen- Risiko-Rang (hoch- niedrig)	Aufteilung s-Rang (hoch- niedrig)
5												
6		Asset Class 1	10.54%	12.36%	10.00%	5.00%	35.00%	0.8524	9	2	7	1
7		Asset Class 2	11.25%	16.23%	10.00%	5.00%	35.00%	0.6929	7	8	10	1
8		Asset Class 3	11.84%	15.64%	10.00%	5.00%	35.00%	0.7570	6	7	9	1
9		Asset Class 4	10.64%	12.35%	10.00%	5.00%	35.00%	0.8615	8	1	5	1
10		Asset Class 5	13.25%	13.28%	10.00%	5.00%	35.00%	0.9977	5	4	2	1
11		Asset Class 6	14.21%	14.39%	10.00%	5.00%	35.00%	0.9875	3	6	3	1
12		Asset Class 7	15.53%	14.25%	10.00%	5.00%	35.00%	1.0898	1	5	1	1
13		Asset Class 8	14.95%	16.44%	10.00%	5.00%	35.00%	0.9094	2	9	4	1
14		Asset Class 9	14.16%	16.50%	10.00%	5.00%	35.00%	0.8584	4	10	6	1
15		Asset Class 10	10.06%	12.50%	10.00%	5.00%	35.00%	0.8045	10	3	8	1
16												
17		Portfolio insgesamt	12.6419%	4.58%	100.00%							
18		Renditen-Risiko-Verhältnis	2.7596									
19												
20												
21		Spezifikationen des Optimierungsmodells:										
22												
23												
24		Ziel:	Renditen-Risiko-Verhältnis maximieren (C18)									
25		Entscheidungsvariablen:	Aufteilungsgewichte (E6:E15)									
26		Beschränkungen für die Entscheidungsvariablen:	Erforderliches Minimum und Maximum (F6:G15)									
27		Einschränkungen:	Gesamtaufteilungsgewichte des Portfolios 100% (E17 ist auf 100% eingestellt)									
28												
29		Zusätzliche Spezifikationen:										
30												
31		1. Man kann immer die Gesamtportfoliorenditen maximieren oder das Gesamtportfoliorisiko minimieren.										
32		2. Integrieren Sie eine Monte-Carlo-Simulation ins Modell, indem Sie die Renditen und Volatilität jeder Aktivaklasse simulieren und Simulations-Optimierungsmethoden anwenden.										
33		3. Man kann das Portfolio so wie es ist ohne Simulation optimieren unter Verwendung von statischen Optimierungsmethoden.										

Bild 4.1 - Kontinuierliches Optimierungsmodell

Die Aufteilungsgewichte in Spalte E enthalten die Entscheidungsvariablen, welches die Variablen sind, die man justieren und testen muss, sodass das Gesamtgewicht auf 100% beschränkt ist (Zelle E17). Normalerweise, um die Optimierung zu starten stellt man die Zellen auf einen Uniformwert ein, wobei in diesem Fall die Zellen E6 bis E15 jede auf 10% eingestellt ist.

Außerdem könnte jede Entscheidungsvariable spezifische Beschränkungen in ihrem erlaubten Bereich haben. In diesem Beispiel sind die erlaubten unteren und oberen Aufteilungen 5% und 35%, wie in Spalten F und G angezeigt. Diese Einstellung bedeutet, dass jede Aktivaklasse seine eignen Aufteilungsgrenzen haben könnte. Als nächstes zeigt die Spalte H das Rendite-Risiko-Verhältnis, was einfach der Renditeprozentsatz geteilt durch den Risikoprozentsatz ist, wobei je höher dieser Wert, desto höher das „Bang-for-the-Buck“. Das verbleibende Modell zeigt die Ranglisten der individuellen Aktivaklassen nach Renditen, Risiko, Rendite-Risiko-Verhältnis und Aufteilung. In anderen Worten, diese Ranglisten zeigen auf einen Blick, welche Aktivaklasse das niedrigste Risiko oder die höchste Rendite hat und so weiter.

Die Gesamtrenditen des Portfolios in Zelle C17 ist $SUMPRODUCT(C6:C15, E6:E15)$, das heißt, die Summe der Aufteilungsgewichte multipliziert mit den annualisierten Renditen jeder Aktivaklasse. In anderen Worten, wir haben $R_p = \omega_A R_A + \omega_B R_B + \omega_C R_C + \omega_D R_D$, wobei R_p die Rendite des Portfolios ist, $R_{A,B,C,D}$ die individuellen Renditen der Projekte sind und $\omega_{A,B,C,D}$ die jeweiligen Gewichte oder jeweilige Kapitalaufteilung durch jedes Projekt sind.

Des Weiteren, die Risikodiversifikation des Portfolios in Zelle D17 wird folgendermaßen

berechnet: $\sigma_p = \sqrt{\sum_{i=1}^i \omega_i^2 \sigma_i^2 + \sum_{i=1}^n \sum_{j=1}^m 2\omega_i \omega_j \rho_{i,j} \sigma_i \sigma_j}$. Hier sind ρ_{ij} die jeweiligen

Kreuzkorrelationen zwischen den Aktivaklassen—deshalb, wenn die Kreuzkorrelationen negativ sind, gibt es Risikodiversifikationseffekte und das Portfoliorisiko sinkt. Allerdings, um die Berechnungen hier zu vereinfachen, nehmen wir an, dass es keine Korrelationen zwischen den Aktivaklassen in dieser Portfoliorisikoberechnung gibt, aber wir nehmen dagegen an, dass Korrelationen vorhanden sind, wenn wir eine Simulation auf diese Erträge anwenden, wie wir später sehen werden. Deshalb, anstatt statische Korrelationen zwischen diesen verschiedenen Aktivaklassen anzuwenden, wenden wir die Korrelationen in den Simulationshypothesen selber an, und kreieren dabei ein dynamischeres Verhältnis zwischen den simulierten Renditenwerten.

Letztlich, es wird das Renditen-Risiko-Verhältnis oder Sharpe-Ratio für das Portfolio berechnet. Dieser Wert wird in Zelle C18 angezeigt und repräsentiert das Ziel, das in dieser Optimierungsübung maximiert sein soll. Um zusammenzufassen, wir haben die folgenden Spezifikationen in diesem Beispielsmodell:

Ziel:	<i>Renditen-Risiko-Verhältnis maximieren (C18)</i>
Entscheidungsvariablen:	<i>Aufteilungsgewichte (E6:E15)</i>
Beschränkungen für die Entscheidungsvariablen:	<i>Erforderliches Minimum und Maximum (F6:G15)</i>
Einschränkungen:	<i>Gesamtsumme der Aufteilungsgewichte bis 100% (E17)</i>

Prozedur:

- Die Beispielsdatei öffnen und eine neues Profil starten: Klicken Sie auf **Risiko Simulator** | **Neues Profil** und benennen Sie es.
- Der erste Schritt bei einer Optimierung ist, die Entscheidungsvariablen einzustellen. Wählen Sie die Zelle E6, stellen Sie die erste Entscheidungsvariable ein (**Risiko Simulator** | **Optimierung** | **Entscheidung einstellen**) und klicken Sie auf die Ikone **Verknüpfung**, um den Zellennamen (B6) sowie auch die unteren und oberen Grenzwerte in Zellen F6 und G6 auszuwählen. Dann, unter Verwendung der Option **Kopieren** von Risiko Simulator, kopieren Sie diese Entscheidungsvariable der Zelle E6 und fügen Sie diese Entscheidungsvariable in den restlichen Zellen E7 bis E15 ein.
- Der zweite Schritt bei einer Optimierung ist, die Einschränkung einzustellen. Hier gibt es nur eine Einschränkung, das heißt, die Gesamtaufteilung im Portfolio muss sich auf 100% summieren. Klicken Sie also auf **Risiko Simulator** | **Optimierung** | **Einschränkungen...** und wählen Sie **HINZUFÜGEN**, um eine neue Einschränkung hinzuzufügen. Dann die Zelle E17 wählen und diese gleich (=) 100% einstellen. Wenn fertig, auf OK klicken.
- Der letzte Schritt bei einer Optimierung ist, die Zielfunktion einzustellen und die Optimierung zu starten: Wählen Sie erst die Zielzelle C18, dann **Risiko Simulator** | **Optimierung** | **Optimierung ausführen** und letztlich die gewünschte Optimierung aus (Statische Optimierung, Dynamische Optimierung oder Stochastische Optimierung). Um zu beginnen, wählen Sie **Statische Optimierung**. Vergewissern Sie sich, dass die Zielzelle für C18 eingestellt ist und wählen Sie **Maximieren**. Sie können jetzt die Entscheidungsvariablen und Einschränkungen falls erwünscht prüfen, oder auf OK klicken, um die statische Optimierung auszuführen.
- Wenn die Optimierung abgeschlossen ist, können Sie **Zurückkehren** wählen, um zu den originalen Werten der Entscheidungsvariablen und des Ziels zurückzukehren, oder auf **Ersetzen** klicken, um die optimierten Entscheidungsvariablen anzuwenden. Typischerweise wird **Ersetzen** nach einer abgeschlossenen Optimierung gewählt.

Bild 4.2 zeigt die Bildschirmansichten dieser obigen Verfahrensschritte. Sie können Simulationshypothesen zu den Renditen und Risiken (Spalten C und D) des Modells hinzufügen und die dynamische Optimierung und stochastische Optimierung für zusätzliche Übung anwenden.

Eigenschaften der Entscheidungsvariable

Entscheidungsname: Asset Class 1

Entscheidungstyp:

- ☒ Kontinuierlich (z. B., 1.15, 2.35, 10.55)

Untere Grenze: 0.05 Obere Grenze: 0.35
- ☐ Ganzzahl (z. B., 1, 2, 3)

Untere Grenze: Obere Grenze:
- ☐ Binär (0 oder 1)

OK Abbrechen

Einschränkung

Zelle: \$E\$17 == Einschränkung: 100%

OK Abbrechen

Optimierungszusammenfassung

Die Optimierung wird verwendet, um Ressourcen zu verteilen, wobei die Ergebnisse die Maximalerträge oder die minimalen Kosten/Risiken liefern. Anwendungen sind, unter anderem, Verwaltung von Inventaren, Finanzportfolioallokation, Produktmix, Projektauswahl, usw.

Optimization

Ziel Methode **Einschränkungen** Statistiken Entscheidungsvariablen

- ☒ **Statische Optimierung**

In einem statischen Modell ohne Simulationen ausführen. Normalerweise ausgeführt, um das optimale Anfangsportfolio zu bestimmen, bevor man fortgeschrittenere Optimierungen anwendet.
- ☐ **Dynamische Optimierung**

Erst wird eine Simulation ausgeführt, die Simulationsergebnisse werden im Modell angewendet und dann wird eine Optimierung auf die simulierten Werte angewendet.

Anzahl der Simulationsprobeversuche: 500
- ☐ **Stochastische Optimierung**

Ähnlich der dynamischen Optimierung aber der Prozess wird mehrere Male wiederholt. Die endgültigen Entscheidungsvariablen werden jede ihr eigenes Vorausberechnungsdiagramm aufweisen, die ihren optimalen Bereich anzeigt.

Anzahl der Simulationsprobeversuche: 500

Anzahl der Optimierungsdurchführungen: 20

Fortgeschrittene OK Abbrechen

Bild 4.2 - Eine kontinuierliche Optimierung in Risiko Simulator ausführen

Interpretierung der Ergebnisse:

Die Endergebnisse der Optimierung werden im Bild 4.3 angezeigt, wobei die optimale Aufteilung der Aktiva für das Portfolio in Zellen E6:E15 zu sehen ist. Das heißt, gegeben die Beschränkungen, dass jedes Aktivum zwischen 5% und 35% fluktuieren kann und dass die Summe der Aufteilung gleich 100% sein muss, zeigt das Bild 4.3 die Aufteilung, die das Rendite-Risiko-Verhältnis maximiert.

Man sollte einige wichtige Punkte beachten, wenn man die Ergebnisse und die bisher ausgeführten Optimierungsverfahren überprüft:

- Die angemessene Weise eine Optimierung auszuführen ist, das „Bang-for-the-Buck“ oder Renditen-Risiko-Sharpe-Ratio zu maximieren, sowie wir es getan haben.
- Wenn wir stattdessen die Gesamtportfolioerträge maximieren, ist das optimale Aufteilungsergebnis belanglos und benötigt keine Optimierung, um es zu erhalten. In anderen Worten, teilen Sie einfach 5% (das erlaubte Minimum) zu den niedrigsten 8 Aktiva, 35% (das erlaubte Maximum) zum Aktivum mit der höchsten Rendite und den Rest (25%) zum Aktivum mit der zweithöchsten Rendite auf. Eine Optimierung ist nicht erforderlich. Wenn Sie allerdings das Portfolio auf diese Weise aufteilen, ist das Risiko sehr viel höher im Vergleich zur Maximierung des Renditen-Risiko-Verhältnisses, obwohl die Portfolioerträge selber höher sind.
- Im Gegensatz dazu kann man das Gesamtportfoliorisiko minimieren, aber jetzt werden die Renditen geringer sein.

Die Tabelle 4.1 stellt die Ergebnisse der drei verschiedenen zu optimierenden Ziele dar:

Ziel:	Portfolio-Renditen	Portfolio-Risiko	Renditen-Risiko-Verhältnis des Portfolios
Renditen-Risiko-Verhältnis maximieren	12.69%	4.52%	2.8091
Renditen maximieren	13.97%	6.77%	2.0636
Risiko minimieren	12.38%	4.46%	2.7754

Tabelle 4.1 - Optimierungsergebnisse

Wie aus der Tabelle ersichtlich, ist die beste Methode die Maximierung des Renditen-Risiko-Verhältnisses. Das heißt, für denselben Umfang von Risiko liefert diese Aufteilung den höchsten Renditenbetrag. Umgekehrt, für denselben Renditenbetrag liefert diese Aufteilung den möglichst geringen Umfang von Risiko. Diese Methode des „Bang-for-the-Buck“ oder Renditen-Risiko-Verhältnisses ist der Grundstein der effizienten Grenze von Markowitz in der modernen Portfoliotheorie. In anderen Worten, wenn wir die Gesamtportfoliorisikoniveaus beschränken und diese im Laufe der Zeit stufenweise erhöhen, erhalten wir diverse effiziente Portfolioaufteilungen

für unterschiedliche Risikoeigenschaften. Infolgedessen kann man unterschiedliche effiziente Portfolioaufteilungen für unterschiedliche Individuen mit ungleichen Risikopräferenzen erhalten.

OPTIMIERUNGSMODELL FÜR DIE AUFTEILUNG DER AKTIVA

Beschreibung der Aktivaklasse	Annualisierte Renditen	Volatilitäts-risiko	Aufteilungs-gewichte	Erforderliche minimale Aufteilung	Erforderliche maximale Aufteilung	Renditen-Risiko-Verhältnis
Asset Class 1	10.54%	12.36%	11.09%	5.00%	35.00%	0.8524
Asset Class 2	11.25%	16.23%	6.86%	5.00%	35.00%	0.6929
Asset Class 3	11.84%	15.64%	7.78%	5.00%	35.00%	0.7570
Asset Class 4	10.64%	12.35%	11.23%	5.00%	35.00%	0.8615
Asset Class 5	13.25%	13.28%	12.09%	5.00%	35.00%	0.9977
Asset Class 6	14.21%	14.39%	11.04%	5.00%	35.00%	0.9875
Asset Class 7	15.53%	14.25%	12.30%	5.00%	35.00%	1.0898
Asset Class 8	14.95%	16.44%	8.90%	5.00%	35.00%	0.9094
Asset Class 9	14.16%	16.50%	8.37%	5.00%	35.00%	0.8584
Asset Class 10	10.06%	12.50%	10.35%	5.00%	35.00%	0.8045
Portfolio insgesamt	12.6919%	4.52%	100.00%			
Renditen-Risiko-Verhältnis	2.8091					

Bild 4.3 - Ergebnisse einer kontinuierlichen Optimierung

Optimierung mit diskreten ganzzahligen Variablen

Gelegentlich sind die Entscheidungsvariablen nicht kontinuierlich sondern diskrete Ganzzahlen (z.B., 0 und 1). Das heißt, wir können solch eine Optimierung als Ein-Aus Schalter oder Go/No-Go Entscheidungen verwenden. Bild 4.4 zeigt ein Projektauswahlmodell, wobei 12 Projekte aufgelistet sind. Dieses Beispiel verwendet die Datei **Diskrete Optimierung**, die entweder im Startmenü unter **Start | Real Options Valuation | Risiko Simulator | Beispiele** oder direkt unter **Risiko Simulator | Beispielsmodelle** zu finden ist. Wie zuvor hat jedes Projekt seine eigene Renditen (ENPV für erweiterter Nettogegenwartswert und NPV Nettogegenwartswert – ENPV ist lediglich NPV plus jegliche strategische Realloptionswerte), Implementierungskosten, Risiken und so weiter. Bei Bedarf kann man dieses Modell modifizieren, um erforderliche Vollzeitäquivalenten (FTE) und andere Ressourcen von verschiedenen Funktionen einzuschließen und zusätzliche Einschränkungen auf diese zusätzlichen Ressourcen einstellen. Die Inputs in diesem Modell werden typischerweise von anderen Arbeitsblattmodellen verknüpft. Zum Beispiel, jedes Projekt wird sein eigenes „diskontierter Cashflow“ Modell oder Kapitalrendite-Modell haben. Die Anwendung hier ist die Sharpe-Ratio des Portfolios zu maximieren, abhängig von bestimmten Budgetaufteilungen. Man kann viele andere Versionen dieses Modells kreieren wie, zum Beispiel, die Maximierung der Portfoliorenditen oder die Minimierung der Risiken, oder zusätzliche Einschränkungen hinzufügen, wobei die Gesamtzahl der ausgewählten Projekte 6 nicht überschreiten darf, und so weiter und so fort. Man kann alle diese Elemente unter Verwendung dieses existierenden Modells ausführen.

Prozedur:

- Die Beispielsdatei öffnen und eine neues profil starten: Klicken Sie auf **Risiko Simulator** | **Neues Profil** und benennen Sie es.
- Der erste Schritt bei einer Optimierung ist, die Entscheidungsvariablen einzustellen. Stellen Sie die erste Entscheidungsvariable durch die Auswahl der Zelle J4 ein, wählen Sie **Risiko Simulator** | **Optimierung** | **Entscheidung einstellen**, klicken Sie auf die Ikone **Verknüpfung**, um den Zellennamen (B4) auszuwählen, und wählen Sie dann die Variable **Binär**. Dann, unter Verwendung der Option **Kopieren** von Risiko Simulator, kopieren Sie diese Entscheidungsvariable der Zelle J4 und fügen Sie diese Entscheidungsvariable in die restlichen Zellen J5 bis J15 ein. Dies ist die beste Methode, wenn Sie nur einige Entscheidungsvariablen haben und Sie jede Entscheidungsvariable mit einem einmaligen Namen für spätere Identifizierung benennen können.
- Der erste Schritt bei einer Optimierung ist, die Einschränkung einzustellen. Hier gibt es zwei Einschränkungen, sprich, die Gesamtbudgetaufteilung im Portfolio muss weniger als \$5000 sein und die Gesamtzahl der Projekte darf 6 nicht überschreiten. Klicken Sie also auf **Risiko Simulator** | **Optimierung** | **Einschränkungen...** und wählen Sie **HINZUFÜGEN**, um eine neue Einschränkung hinzuzufügen. Dann die Zelle D17 wählen und diese kleiner als oder gleich ($\leq$) 5000 einstellen. Dies wiederholen, indem Sie die Zelle J17 $\leq$ 6 einstellen.
- Der letzte Schritt bei einer Optimierung ist, die Zielfunktion einzustellen und die Optimierung zu starten: Wählen Sie erst die Zielzelle C19 und dann **Risk Simulator** | **Optimierung** | **Ziel einstellen**. Dann führen Sie die Optimierung unter Verwendung von **Risiko Simulator** | **Optimierung** | **Optimierung ausführen** und der Auswahl der gewünschten Optimierung (Statische Optimierung, Dynamische Optimierung oder Stochastische Optimierung) aus. Um zu beginnen, wählen Sie **Statische Optimierung**. Vergewissern Sie sich, dass die Zielzelle entweder eine Sharpe-Ratio oder ein Renditen-Risiko-Verhältnis ist und wählen Sie **Maximieren**. Sie können jetzt die Entscheidungsvariablen und Einschränkungen falls erwünscht prüfen, oder auf OK klicken, um die statische Optimierung auszuführen.

Bild 4.5 zeigt die Bildschirmansichten dieser obigen Verfahrensschritte. Sie können Simulationshypothesen zu den erweiterten Nettogegenwartswert (ENPV) und Risiken (Spalten C und D) des Modells hinzufügen und die dynamische Optimierung und stochastische Optimierung für zusätzliche Übung anwenden.

	A	B	C	D	E	F	G	H	I	J	K
1											
2											
3		Projekte	ENPV	Kosten	Risiko \$	Risiko %	Rendite- Risiko- Verhältnis	Profitabilitäts- index		Auswahl	
4		Projekt 1	\$458.00	\$1,732.44	\$54.96	12.00%	8.33	1.26		1.0000	
5		Projekt 2	\$1,954.00	\$859.00	\$1,914.92	98.00%	1.02	3.27		1.0000	
6		Projekt 3	\$1,599.00	\$1,845.00	\$1,551.03	97.00%	1.03	1.87		1.0000	
7		Projekt 4	\$2,251.00	\$1,645.00	\$1,012.95	45.00%	2.22	2.37		1.0000	
8		Projekt 5	\$849.00	\$458.00	\$925.41	109.00%	0.92	2.85		1.0000	
9		Projekt 6	\$758.00	\$52.00	\$560.92	74.00%	1.35	15.58		1.0000	
10		Projekt 7	\$2,845.00	\$758.00	\$5,633.10	198.00%	0.51	4.75		1.0000	
11		Projekt 8	\$1,235.00	\$115.00	\$926.25	75.00%	1.33	11.74		1.0000	
12		Projekt 9	\$1,945.00	\$125.00	\$2,100.60	108.00%	0.93	16.56		1.0000	
13		Projekt 10	\$2,250.00	\$458.00	\$1,912.50	85.00%	1.18	5.91		1.0000	
14		Projekt 11	\$549.00	\$45.00	\$263.52	48.00%	2.08	13.20		1.0000	
15		Projekt 12	\$525.00	\$105.00	\$309.75	59.00%	1.69	6.00		1.0000	
16											
17		Summe	\$17,218.00	\$8,197.44	\$7,007	40.70%				12.00	
18		Ziel:	MAX	<=\$5000						<=6	
19		Sharpe-Ratio	2.4573								
20											
21		ENPV ist der Erwartungsnettogegenwartswert jeder Investition oder jedes Projekts, während Kosten die Gesamtkosten der									
22		Investition sein könnte und Risiko ist der Variationskoeffizient des Erwartungsnettogegenwartswertes (ENPV) des Projekts.									

Bild 4.4 - Diskrete Ganzzahl-Optimierung

Eigenschaften der Entscheidungsvariable

Entscheidungsname

Entscheidungstyp

☐ Kontinuierlich (z.B., 1.15, 2.35, 10.55)

Untere Grenze
Obere Grenze

☐ Ganzzahl (z.B., 1, 2, 3)

Untere Grenze
Obere Grenze

☒ Binär (0 oder 1)

OK

Abbrechen

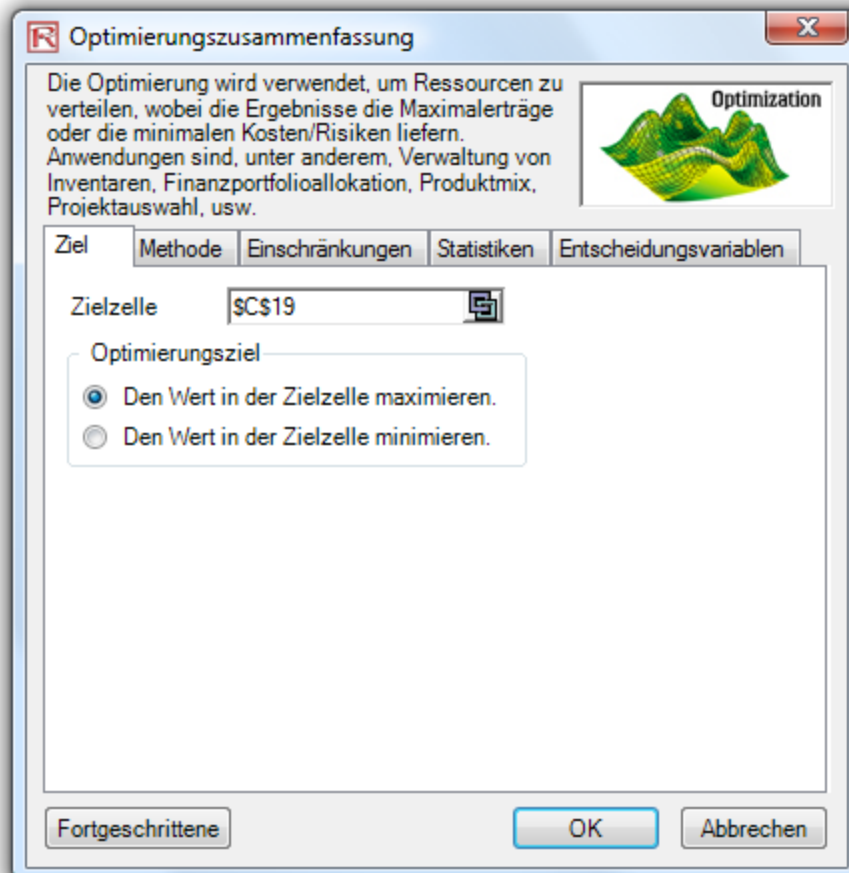
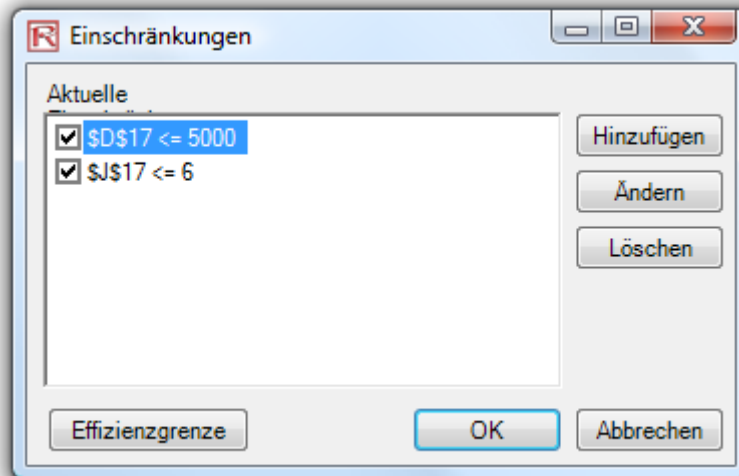


Bild 4.5 - Eine diskrete Ganzzahl-Optimierung in Risiko Simulator ausführen

Interpretierung der Ergebnisse:

Bild 4.6 zeigt das Beispiel einer optimalen Auswahl von Projekten, welche die Sharpe-Ratio maximiert. Im Gegensatz dazu kann man immer die Gesamtgewinne maximieren, aber wie zuvor, ist das ein belangloser Prozess und impliziert lediglich die Auswahl des Projekts mit der höchsten Rendite und das Heruntergehen der Liste, bis die Mittel ausgehen oder die Budgeteinschränkung überschritten wird. Diese Verfahrensweise könnte theoretische nicht wünschenswerte Projekte ergeben, da die höchstergebenden Projekte typisch die höchsten Risiken besitzen. Jetzt, wenn erwünscht, können Sie die Optimierung unter Verwendung einer stochastischen oder dynamischen Optimierung replizieren, indem Sie Hypothesen in den Werten für ENPV und/oder Kosten und/oder Risiko hinzufügen.

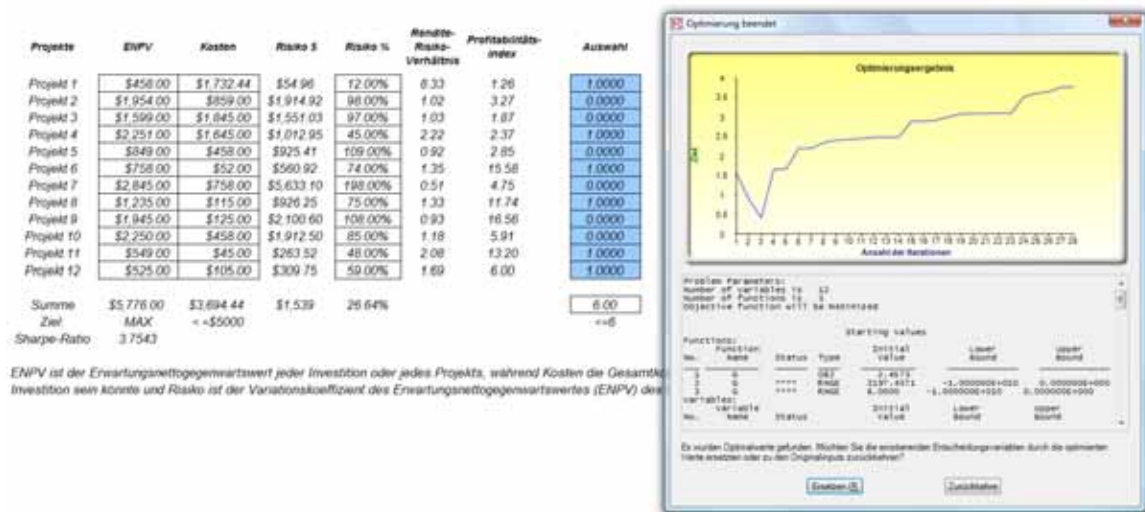


Bild 4.6 - Optimale Auswahl von Projekten, welche die Sharpe-Ratio maximieren

Für zusätzliche praktische Beispiele der Optimierung, siehe die Fallstudie über *Integrated Risk Analysis* im Kapitel 11 des Buchs *Real Options Analysis: Tools and Techniques*, 2nd Edition (Wiley Finance, 2005). Diese Fallstudie erläutert wie man eine effiziente Grenze generieren und wie man Vorausberechnung, Simulation, Optimierung und Realoptionen in einem nahtlosen analytischen Prozess kombinieren kann.

Effiziente Grenze und fortgeschrittene Optimierungseinstellungen

Das zweite Diagramm im Bild 4.5 zeigt die Einschränkungen für die Optimierung. Wenn Sie auf die Schaltfläche **Effiziente Grenze nach** der Einstellung einiger Einschränkungen geklickt haben, können Sie hier diese Einschränkungen als wechselnd einstellen. In anderen Worten, man kann jede der Einschränkungen so kreieren, dass diese schrittweise zwischen einem bestimmten Maximal- und Minimalwert durchgeführt werden. Als Beispiel, man kann die Einschränkung in Zelle J17 ≤ 6 so einstellen, dass sie zwischen 4 und 8 durchgeführt wird (Bild 4.7). Das heißt, es werden fünf Optimierungen ausgeführt, jede mit den folgenden Einschränkungen: J17 ≤ 4 , J17 ≤ 5 , J17 ≤ 6 , J17 ≤ 7 und J17 ≤ 8 . Die optimalen Ergebnisse werden dann als eine effiziente Grenze geplottet und der Bericht wird generiert (Bild 4.8). Im Besonderen erläutert folgendes die erforderlichen Schritte, um eine wechselnde Einschränkung zu kreieren:

- ❑ In ein Optimierungsmodell (das heißt, ein Modell mit bereits eingestelltem Ziel, Entscheidungsvariablen und Einschränkungen), klicken Sie auf **Risiko Simulator | Optimierung | Einschränkungen** und dann auf **Effiziente Grenze**.
- ❑ Wählen Sie die zu wechselnde oder schrittweise auszuführende Einschränkung (z.B., J17), geben Sie die Parameter für Min, Max und Schrittgröße ein (Bild 4.7) und klicken Sie auf **HINZUFÜGEN** und dann auf **OK** und erneut auf **OK**. Sie sollten die Einschränkung D17 ≤ 5000 vor der Ausführung deaktivieren.
- ❑ Führen Sie die Optimierung wie gewohnt aus (**Risiko Simulator | Optimierung | Optimierung ausführen**). Sie können statische, dynamische oder stochastische auswählen.
- ❑ Die Ergebnisse werden als eine Benutzeroberfläche angezeigt (Bild 4.8). Klicken Sie auf **Bericht kreieren**, um ein Berichtsarbeitsblatt mit allen Details der Optimierungsläufe zu generieren.

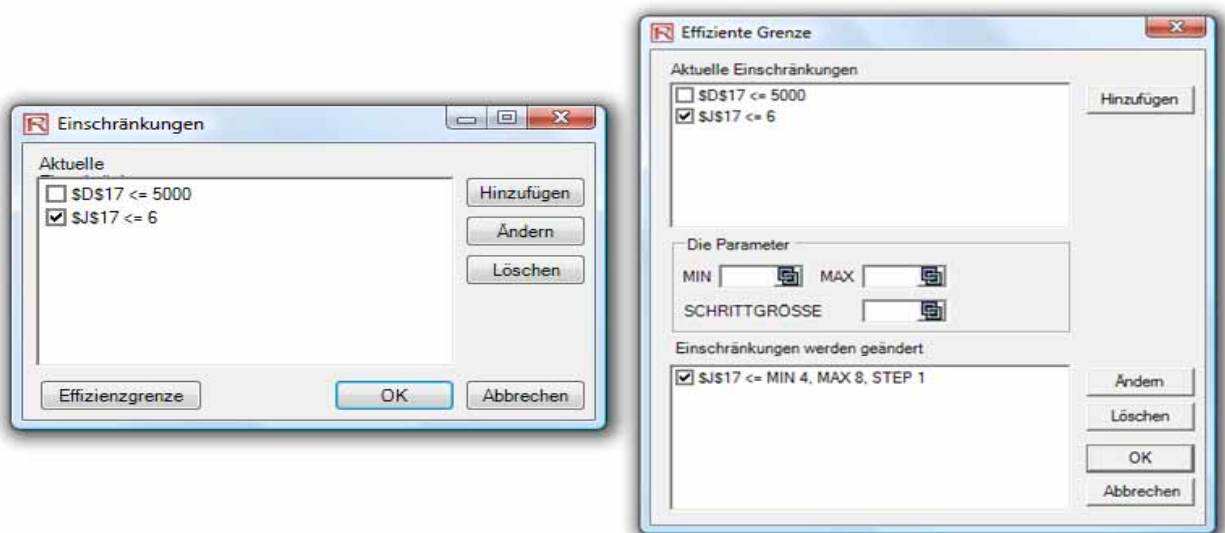
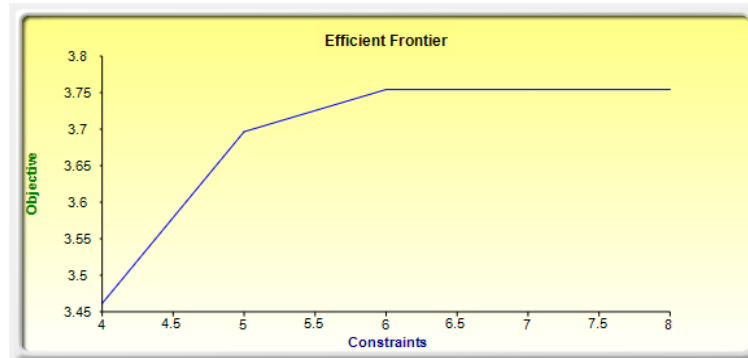


Bild 4.7 - Wechselnde Einschränkungen in einer effizienten Grenze generieren

Effiziente Grenze

Problem Parameters:
 Number of variables 12
 Number of functions 2
 Objective function will be Maximized



STEP1, D17 <= 5000, J17 <= 4

Functions

Starting Values							Final Results					
No.	Function Name	Status	Type	Initial Value	Lower Bound	Upper Bound	No.	Function Name	Initial Value	Final Value	Status	Distance from Nearest Bound
1	G		OBJ	2.45726			1	G	2.45726	3.46137	Objective	
2	G	****	RNGE	8.00000	-1E+10	0	2	G	8.00000	0.00000	UpperBnd	0.0000:U 0.23768

Variables

Starting Values						Final Results					
No.	Variable Name	Status	Initial Value	Lower Bound	Upper Bound	No.	Variable Name	Initial Value	Final Value	Status	Distance from Nearest Bound
1	Proiekt 1	UL	1.00000	0	1	1	Proiekt 1	1.00000	1.00000	NonBasic	UpperBnd 0.0511

Bild 4.8 - Ergebnisse der effizienten Grenze

Stochastische Optimierung

Das nächste Beispiel erläutert die Anwendung einer stochastischen Optimierung unter Verwendung eines Beispielsmodells mit vier Aktivaklassen, jede mit unterschiedlichen Risiko- und Renditeigenschaften. Die Idee hier ist, die beste Portfolioaufteilung zu finden, sodass das „Bang-for-the-Buck“ oder Renditen-Risiko-Verhältnis des Portfolios maximiert wird. Das heißt, das Ziel ist die Aufteilung von 100% der Investitionen eines Individuums unter mehreren verschiedenen Aktivaklassen (z.B., verschiedene Typen von Anlagefonds oder Investitionsarten: Wachstum, Wert, aggressives Wachstum, Einkommen, global, Index, nonkonformistisch, Momentum und so weiter). Dieses Modell unterscheidet sich von anderen, weil einige Simulationshypothesen (Risiko- und Renditenwerte für jedes Aktivun in Spalten C und D) vorhanden sind, wie im Bild 4.9 angezeigt.

Erst wird eine Simulation ausgeführt, dann wird eine Optimierung durchgeführt und das gesamte Verfahren wird mehrfach wiederholt, um Verteilungen für jede Entscheidungsvariable zu

erhalten. Man kann die gesamte Analyse unter Verwendung der stochastischen Optimierung automatisieren.

Um die Optimierung auszuführen, muss man erst einige Hauptspezifikationen des Modells identifizieren:

Ziel: Renditen-Risiko-Verhältnis maximieren (C12)

Entscheidungsvariablen: Aufteilungsgewichte (E6:E9)

Beschränkungen für die Entscheidungsvariablen: Erforderliches Minimum und Maximum (F6:G9)

Einschränkungen: Gesamtaufteilungsgewichte des Portfolios 100% (E11 ist auf 100% eingestellt)

Simulationshypothesen: Renditen- und Risikowerte (C6:D9)

Das Modell zeigt die verschiedenen Aktivaklassen. Jede Aktivaklasse hat seinen eigenen Satz von annualisierten Renditen und annualisierten Volatilitäten. Diese Renditen- und Risikomessungen sind annualisierte Werte, sodass man sie konsequent über verschiedene Aktivaklassen vergleichen kann. Renditen werden unter Verwendung des geometrischen Mittels der relativen Renditen berechnet, während die Risiken unter Verwendung der Methode der logarithmischen relativen Aktienrenditen berechnet werden.

Die Aufteilungsgewichte in Spalte E enthalten die Entscheidungsvariablen, welche die Variablen sind, die man justieren und testen muss, sodass das Gesamtgewicht auf 100% beschränkt ist (Zelle E11). Normalerweise, um die Optimierung zu starten, stellt man die Zellen auf einen Uniformwert ein. In diesem Fall sind die Zellen E6 bis E9 jede auf 25% eingestellt. Außerdem könnte jede Entscheidungsvariable spezifische Beschränkungen in ihrem erlaubten Bereich haben. In diesem Beispiel, sind die erlaubten unteren und oberen Aufteilungen 10% und 40%, wie in Spalten F und G angezeigt. Diese Einstellung bedeutet, dass jede Aktivaklasse seine eigenen Aufteilungsgrenzen haben könnte.

	A	B	C	D	E	F	G	H
1								
2								
3								
4								
5								
6								
7								
8								
9								
10								
11								
12								
13								

OPTIMIERUNGSMODELL FÜR DIE AUFTEILUNG DER AKTIVA							
	Beschreibung der Aktivaklasse	Annualisierte Renditen	Volatilitäts-risiko	Aufteilungs-gewichte	Erforderliche minimale Aufteilung	Erforderliche maximale Aufteilung	Renditen-Risiko-Verhältnis
6	Asset 1	10.60%	12.41%	25.00%	10.00%	40.00%	0.8544
7	Asset 2	11.21%	16.16%	25.00%	10.00%	40.00%	0.6937
8	Asset 3	10.61%	15.93%	25.00%	10.00%	40.00%	0.6660
9	Asset 4	10.52%	12.40%	25.00%	10.00%	40.00%	0.8480
11	Portfolio insgesamt	10.7356%	7.17%	100.00%			
12	Renditen-Risiko-Verhältnis	1.4970					

Bild 4.9 - Aktiva Aufteilungsmodell, bereit für eine stochastische Optimierung

Als nächstes zeigt die Spalte H das Renditen-Risiko-Verhältnis, was einfach der Renditenprozentsatz geteilt durch den Risikoprozentsatz für jedes Aktivum ist, wobei je höher dieser Wert, desto höher das „Bang-for-the-Buck“. Die restlichen Teile des Modells zeigen die Ranglisten der individuellen Aktivaklassen nach Renditen, Risiko, Renditen-Risiko-Verhältnis und Aufteilung. In anderen Worten, diese Ranglisten zeigen auf einen Blick, welche Aktivaklasse das niedrigste Risiko oder die höchste Rendite hat und so weiter.

Eine Optimierung ausführen

Um dieses Modell auszuführen, klicken Sie einfach auf **Risiko Simulator | Optimierung | Optimierung ausführen**. Alternativ, zur Übung, können Sie das Modell unter Verwendung der folgenden Schritte einstellen.

1. Starten Sie ein neues Profil (**Risiko Simulator | Neues Profil**).
2. Für eine stochastische Optimierung, stellen Sie die Verteilungshypothesen für das Risiko und die Renditen jeder Aktivaklasse ein. Das heißt, wählen Sie die Zelle **C6** aus und stellen Sie eine Hypothese ein (**Risiko Simulator | Inputhypothese einstellen**) und erstellen Sie bei Bedarf Ihre eigene Hypothese. Wiederholen Sie den Vorgang für Zellen **C7** bis **D9**.
3. Wählen Sie die Zelle **E6** aus und definieren Sie die Entscheidungsvariable (**Risiko Simulator | Optimierung | Entscheidung einstellen** oder klicken Sie auf die Ikone *Entscheidung einstellen D*) und legen Sie diese als eine **kontinuierliche Variable** fest. Dann verknüpfen Sie den Namen der Entscheidungsvariablen und das erforderliche Minimum/Maximum mit den relevanten Zellen (**B6, F6, G6**).
4. Dann verwenden Sie die Funktion **Kopieren** von **Risiko Simulator** auf der Zelle **E6**, wählen Sie die Zellen **E7 bis E9** und verwenden Sie die Funktion **Einfügen** von **Risiko Simulator** (**Risiko Simulator | Parameter kopieren** und **Risiko Simulator | Parameter einfügen** oder verwenden Sie die Ikonen Kopieren und Einfügen). Beachten Sie immer, nicht die normalen Funktionen Kopieren und Einfügen von Excel zu verwenden.
5. Als nächstes stellen Sie die Optimierungseinschränkungen auf: wählen Sie erst **Risiko Simulator | Optimierung | Einschränkungen**, dann **HINZUFÜGEN**, dann wählen Sie die Zelle **E11** aus und stellen Sie diese auf **100%** (Gesamtaufteilung und vergessen Sie nicht das % Symbol einzugeben).
6. Wählen Sie die Zelle **C12**, das zu maximierende Ziel und machen Sie es das Ziel: **Risiko Simulator | Optimierung | Ziel einstellen** oder klicken Sie auf die Ikone **O**.
7. Führen Sie die Optimierung aus: Gehen Sie zu **Risiko Simulator | Optimierung | Optimierung ausführen**. Überprüfen Sie die verschiedenen Leisten, um sicherzugehen, dass alle erforderlichen Inputs in Schritten 2 und 3 richtig sind. Wählen Sie **Stochastische Optimierung** und diese so einstellen, dass sie 20 x 500 Probeversuche ausführt (Bild 4.10 zeigt diese Einstellungsschritte).

Eigenschaften der Entscheidungsvariable

Entscheidungsname:

Entscheidungstyp:

- ☒ Kontinuierlich (z.B., 1.15, 2.35, 10.55)

Untere Grenze: Obere Grenze:
- ☐ Ganzzahl (z.B., 1, 2, 3)

Untere Grenze: Obere Grenze:
- ☐ Binär (0 oder 1)

OK Abbrechen

Einschränkungen

Aktuelle

☒ \$E\$11 == 100%

Hinzufügen
Ändern
Löschen

Effizienzgrenze OK Abbrechen

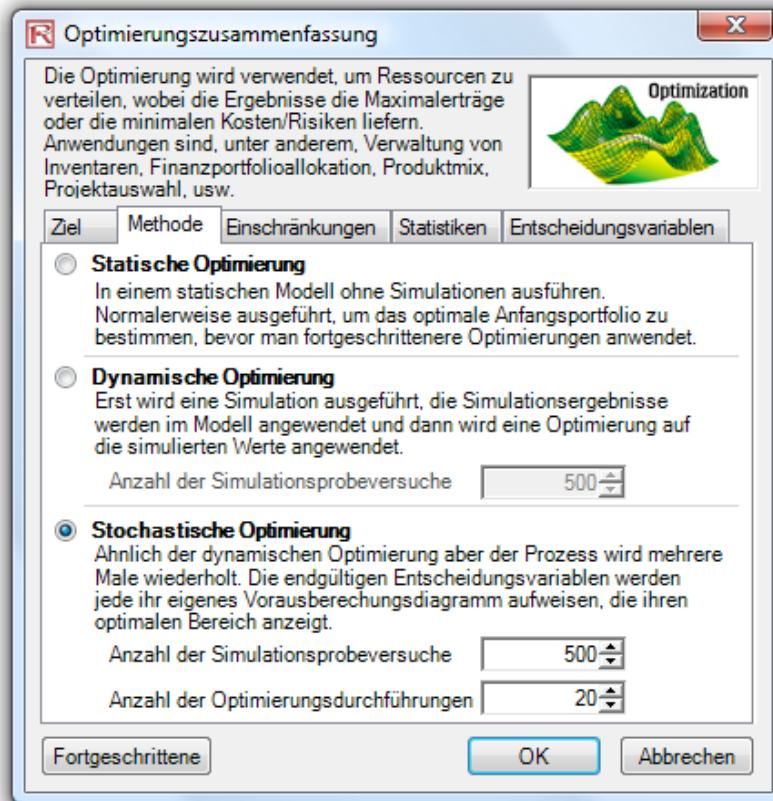


Bild 4.10 - Die stochastische Optimierungsaufgabe

Klicken Sie auf **OK** nach Beendigung der Optimierung. Es werden sowohl ein detaillierter Bericht der stochastischen Optimierung als auch die Vorausberechnungsdiagramme der Entscheidungsvariablen generiert.

Die Vorausberechnungsergebnisse anschauen und interpretieren

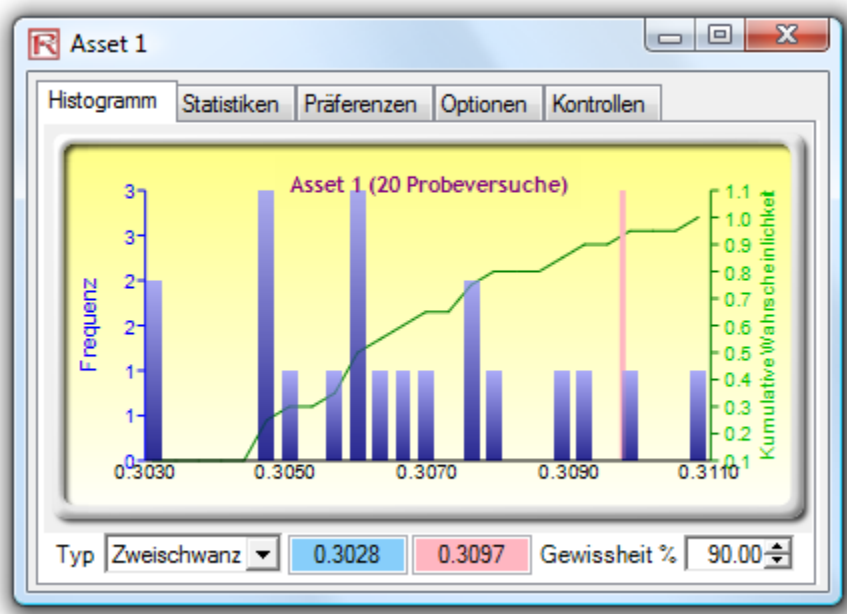
Eine stochastische Optimierung wird ausgeführt, wenn man erst eine Simulation und dann die Optimierung durchführt. Dann wird die gesamte Analyse mehrfach wiederholt. Das Ergebnis ist eine Verteilung jeder Entscheidungsvariablen statt einer Einzelpunktschätzung (Bild 4.11). Das bedeutet, dass anstatt zu behaupten, dass man 30.57% in Aktivum 1 investieren sollte, ist die optimale Entscheidung zwischen 30.10% und 30.99% zu investieren, solange das gesamte Portfolio sich auf 100% summiert. Auf diese Weise liefern die Ergebnisse dem Management oder den Entscheidungsträgern einen Bereich von Flexibilität für die optimalen Entscheidungen, und das unter Berücksichtigung der Risiken und Ungewissheiten in den Inputs.

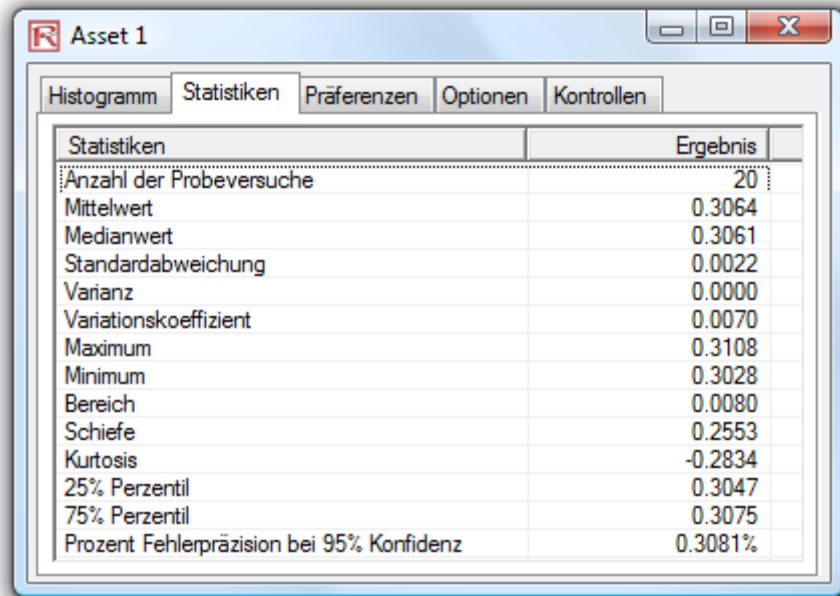
Bemerkungen:

- **Super-Speed-Simulation mit Optimierung.** Man kann auch eine stochastische Optimierung mit Super-Speed-Simulation ausführen. Um so vorzugehen, setzen Sie erst die Optimierung zurück, indem Sie alle vier Entscheidungsvariablen zurück auf 25%

setzen. Dann klicken Sie auf **Optimierung ausführen**, dann auf die Schaltfläche **Fortgeschrittene** (Bild 4.10) und dann aktivieren Sie das Kontrollkästchen für **Super-Speed-Simulation ausführen**. Als nächstes, in der Benutzeroberfläche **Optimierung ausführen** wählen Sie **Stochastische Optimierung** in der Leiste **Methode**, stellen Sie diese für **500** Probeversuche und **20** Optimierungsläufe ein und klicken Sie auf **OK**. Diese Methode kombiniert die Super-Speed-Simulation mit einer Optimierung. Bitte bemerken Sie wie viel schneller die stochastische Optimierung ausgeführt wird. Jetzt können Sie die Optimierung schnell mit einer höheren Anzahl von Simulationsprobeversuchen erneut ausführen.

- **Simulationsstatistiken für stochastische und dynamische Optimierungen.** Bitte bemerken Sie, dass wenn Input-Simulationshypothesen im Optimierungsmodell vorhanden sind (das heißt, diese Inputhypothesen sind erforderlich, um die dynamischen oder stochastischen Optimierungsroutinen auszuführen), ist die Leiste **Statistiken** in der Benutzeroberfläche **Optimierung ausführen** aufgefüllt. Sie können die gewünschten Statistiken aus der Dropdownliste wählen, wie Mittelwert, Standardabweichung, Variationskoeffizient, bedingter Mittelwert, bedingte Varianz, ein spezifisches Perzentil und so weiter. Das bedeutet, dass wenn Sie eine stochastische Optimierung ausführen, wird erst eine Simulation von Tausenden von Probeversuchen durchgeführt, dann wird die ausgewählte Statistik berechnet und dieser Wert wird vorübergehend in die Zelle der Simulationshypothese gestellt. Jetzt wird wieder eine Optimierung auf Basis dieser Statistik ausgeführt und dann wird das Verfahren mehrfach wiederholt. Diese Methode ist wichtig und nützlich für Bankanwendungen bei der Berechnung vom bedingten Wert im Risiko („Conditional VaR“).





Statistiken	Ergebnis
Anzahl der Probeversuche	20
Mittelwert	0.3064
Medianwert	0.3061
Standardabweichung	0.0022
Varianz	0.0000
Variationskoeffizient	0.0070
Maximum	0.3108
Minimum	0.3028
Bereich	0.0080
Schiefe	0.2553
Kurtosis	-0.2834
25% Perzentil	0.3047
75% Perzentil	0.3075
Prozent Fehlerpräzision bei 95% Konfidenz	0.3081%

Bild 4.11 - Simulierte Ergebnisse von der Methode der stochastischen Optimierung

5. ANALYTISCHE TOOLS IN RISIKO SIMULATOR

Dieses Kapitel befasst sich mit den analytischen Tools von Risiko Simulator. Diese analytischen Tools werden unter Verwendung von Beispielsanwendungen der Software Risiko Simulator behandelt, komplett mit Schritt für Schritt Erklärungen. Diese Tools sind sehr wertvoll für Analysten, die im Bereich der Risikoanalyse tätig sind. Die Anwendbarkeit jedes Tools wird im Detail in diesem Kapitel behandelt.

Tornado und Sensibilität Tools in der Simulation

Theorie:

Eines der leistungsstarken Simulationstools ist die Tornadoanalyse—sie erfasst die statischen Auswirkungen jeder Variablen auf das Ergebnis des Modells. Das heißt, das Tool stört automatisch jede Variable im Modell um eine vordefinierte Menge, erfasst die Fluktuation auf der Vorausberechnung oder dem Endergebnis des Modells und listet die resultierenden Störungen von den am meisten zu den am wenigsten signifikanten auf. Bilder 5.1 bis 5.6 erläutern die Anwendung einer Tornadoanalyse. Zum Beispiel, Bild 5.1 ist ein Beispiel eines diskontierten Cashflow Modells, wo die Inputhypothesen des Modells angezeigt werden. Die Frage ist, welche sind die kritischen Erfolgstreiber, die den Output des Modells am meisten beeinflussen? In anderen Worten, was treibt in Wirklichkeit den Nettogegenwartswert von \$96.63 oder welche Inputvariable beeinflusst diesen Wert am meisten?

Man bekommt Zugang zum Tool des Tornadodiagramms unter *Risiko Simulator | Tools | Tornadoanalyse*. Um das erste Beispiel mitzuverfolgen, öffnen Sie die Datei *Tornado- und Sensibilitätsdiagramme (linear)* im Ordner Beispiele. Bild 5.2 zeigt dieses Beispielsmodell, wobei die Zelle G6, die den Nettogegenwartswert enthält, als das zu analysierende Zielergebnis ausgewählt wird. Die im Modell vorausgehenden Variablen („Precedents“) der Zielzelle werden verwendet, um das Tornadodiagramm zu kreieren. „Precedents“ sind alle Input- und Zwischenvariablen, die das Ergebnis des Modells beeinflussen. Zum Beispiel, wenn das Modell aus $A = B + C$ besteht, wobei $C = D + E$, dann sind B, D und E die vorausgehenden Variablen („Precedents“) für A (C ist keine vorausgehende Variable („Precedent“), da sie nur ein berechneter Zwischenwert ist). Bild 5.2 zeigt auch den Testbereich jeder vorausgehenden Variablen, der verwendet wird, um das Zielergebnis zu schätzen. Wenn die vorausgehenden Variablen nur einfache Inputs sind, dann wird der Testbereich eine einfache Störung basierend auf dem ausgewählten Bereich sein (z.B., die Standardeinstellung ist $\pm 10\%$). Man kann bei Bedarf jede vorausgehende Variable bei unterschiedlichen Prozentsätzen stören. Ein breiterer Bereich ist wichtig, weil er besser imstande ist, um für Extremwerte statt für kleinere Störungen um die Erwartungswerten zu testen. Unter bestimmten Umständen können Extremwerte eine größere, kleinere oder unausgeglichene Auswirkung haben (z.B., es können sich Nichtlinearitäten ereignen, da wo steigende oder sinkende Skaleneffekte und Verbundeffekte sich in größere oder

kleinere Werte einer Variablen hereinschleichen) und nur ein breiterer Bereich wird diese nicht linearer Auswirkung erfassen.

	A	B	C	D	E	F	G
1							
2							
3							
4		Basisjahr	2005		Summe PV Nettovorteile		\$1,896.63
5		Markt Risikokorrigierter Diskontsatz	15.00%		Summe PV Investitionen		\$1,800.00
6		Privat-Risiko Diskontsatz	5.00%		Nettgegenwartswert (NPV)		\$96.63
7		Annualisierte Verkaufswachstumsrate	2.00%		Interner Zinsfuß		18.80%
8		Preiserosionsrate	5.00%		Kapitalrendite (ROI)		5.37%
9		Effektiver Steuersatz	40.00%				
10							
11				2005	2006	2007	2008
12		Prod A Durchschnittspreis		\$10.00	\$9.50	\$9.03	\$8.57
13		Prod B Durchschnittspreis		\$12.25	\$11.64	\$11.06	\$10.50
14		Prod C Durchschnittspreis		\$15.15	\$14.39	\$13.67	\$12.99
15		Prod A Menge		50.00	51.00	52.02	53.06
16		Prod B Menge		35.00	35.70	36.41	37.14
17		Prod C Menge		20.00	20.40	20.81	21.22
18		Gesamteinnahmen		\$1,231.75	\$1,193.57	\$1,156.57	\$1,120.71
19		Kosten der verkauften Waren		\$184.76	\$179.03	\$173.48	\$168.11
20		Bruttogewinn		\$1,046.99	\$1,014.53	\$983.08	\$952.60
21		Betriebsausgaben		\$157.50	\$160.65	\$163.86	\$167.14
22		VVG-Kosten		\$15.75	\$16.07	\$16.39	\$16.71
23		Betriebseinkommen (EBITDA)		\$873.74	\$837.82	\$802.83	\$768.75
24		Abschreibung		\$10.00	\$10.00	\$10.00	\$10.00
25		Amortisierung		\$3.00	\$3.00	\$3.00	\$3.00
26		EBIT		\$860.74	\$824.82	\$789.83	\$755.75
27		Zinsausgaben		\$2.00	\$2.00	\$2.00	\$2.00
28		EBT		\$858.74	\$822.82	\$787.83	\$753.75
29		Steuern		\$343.50	\$329.13	\$315.13	\$301.50
30		Nettoeinkommen		\$515.24	\$493.69	\$472.70	\$452.25
31		Abschreibung		\$13.00	\$13.00	\$13.00	\$13.00
32		Änderung im Nettobetriebskapital		\$0.00	\$0.00	\$0.00	\$0.00
33		Kapitalaufwendungen		\$0.00	\$0.00	\$0.00	\$0.00
34		Freier Cashflow		\$528.24	\$506.69	\$485.70	\$465.25
35							
36		Investitionen		\$1,800.00			
37							

Bild 5.1 - Beispielsmodell

Prozedur:

- Wählen Sie die Einzeloutputzelle (das heißt, eine Zelle mit einer Funktion oder Gleichung) in einem Excelmodell (z.B., Zelle G6 wurde in unserem Beispiel ausgewählt)
- Wählen Sie **Risiko Simulator | Tools | Tornadoanalyse**
- Prüfen Sie die vorausgehenden Variablen („Precedents“) und benennen Sie dieses angemessenerweise um (das Umbenennen der vorausgehenden Variablen („Precedents“) auf kürzeren Namen erlaubt ein optisch angenehmeres Tornado- und Spinnennetzdiagramm). Dann klicken Sie auf **OK**

Diskontierter Cashflow Modell					
Basissjahr:	2005	Summe PV Nettovorteile:	\$1,896.63		
Markt Risikokorrigierter Diskontsatz:	15.00%	Summe PV Investitionen:	\$1,800.00		
Privat Risiko Diskontsatz:	5.00%	Nettogegenwartswert (NPV):	\$96.63		
Annualisierte Verkaufswachstumsrate:	2.00%	Interner Zinsfuß:	18.50%		
Preiserlösnormale:	5.00%	Kapitalrendite (ROE):	5.37%		
Effektiver Steuersatz:	40.00%				
	2005	2006	2007	2008	2009
Prod A Durchschnittspreis	\$10.00	\$9.50	\$9.03	\$8.57	\$8.15
Prod B Durchschnittspreis	\$12.25	\$11.64	\$11.06	\$10.50	\$9.98
Prod C Durchschnittspreis	\$15.15	\$14.39	\$13.67	\$12.99	\$12.34
Prod A Menge	50.00	51.00	52.02	53.06	54.12
Prod B Menge	35.00	35.70	36.41	37.14	37.89
Prod C Menge	20.00	20.40	20.81	21.22	21.65
Gesamteinnahmen	\$1,231.75	\$1,193.67	\$1,156.57	\$1,120.71	\$1,085.97
Kosten der verkauften Waren	\$154.76	\$179.03	\$173.48	\$168.11	\$162.00
Bruttogewinn	\$1,046.99	\$1,014.53	\$983.08	\$952.60	\$923.97
Betriebsausgaben	\$157.50	\$160.65	\$163.86	\$167.14	\$170.48
VVG-Kosten	\$15.75	\$16.07	\$16.39	\$16.71	\$17.05
Betriebseinkommen (EBITDA)	\$873.74	\$827.82	\$802.83	\$768.75	\$736.54
Abgeschrieben	\$10.00	\$10.00	\$10.00	\$10.00	\$10.00
Amortisierung	\$3.00	\$3.00	\$3.00	\$3.00	\$3.00
EBIT	\$860.74	\$824.82	\$789.83	\$755.75	\$722.54
Zinsausgaben	\$2.00	\$2.00	\$2.00	\$2.00	\$2.00
EBT	\$858.74	\$822.82	\$787.83	\$753.75	\$720.54
Steuern	\$343.50	\$329.13	\$315.13	\$301.50	\$288.22
Nettoeinkommen	\$515.24	\$493.69	\$472.70	\$452.25	\$432.32
Abgeschrieben	\$13.00	\$13.00	\$13.00	\$13.00	\$13.00
Änderung im Nettobetriebskapital	\$0.00	\$0.00	\$0.00	\$0.00	\$0.00
Kapitalaufwendungen	\$0.00	\$0.00	\$0.00	\$0.00	\$0.00
Freier Cashflow	\$528.24	\$506.69	\$489.70	\$465.25	\$448.32
Investitionen	\$1,800.00				

Bild 5.2—Eine Tornadoanalyse ausführen

Interpretierung der Ergebnisse:

Bild 5.3 präsentiert den resultierenden Bericht der Tornadoanalyse, welcher anzeigt, dass die Kapitalinvestition die größte Auswirkung auf den Nettogegenwartswert hat, gefolgt vom Steuersatz, vom Durchschnittsverkaufspreis und von der verlangten Menge der Produktlinien und so weiter. Der Bericht enthält vier verschiedene Elemente:

- ② Die Statistische Zusammenfassung, welche die ausgeführten Prozeduren auflistet.
- ② Die Sensibilitätstabelle (Bild 5.4) zeigt den anfänglichen NPV Basiswert von 96.63 und wie jedes Input geändert wird (z.B., Investitionen ändert sich von \$1800 auf \$1980 im Vorteil mit einer Spanne von +10%, und von \$1800 auf \$1620 im Nachteil mit einer Spanne von -10% swing). Die resultierende Vorteils- und Nachteilswerte für NPV sind – \$83.37 und \$276.63 mit einer Gesamtänderung von \$360, was diese Variable, die mit der höchsten Auswirkung auf NPV macht. Die vorausgehenden Variablen sind nach der höchsten runter zur geringsten Auswirkung geordnet.
- ② Das Spinnennetzdiagramm (Bild 5.5) stellt diese Effekte graphisch dar. Die y-Achse ist der NPV Zielwert, während die x- Achse, die prozentuale Änderung auf jedem der vorausgehenden Werte ist (der zentraler Punkt ist der Basisfallwert von 96.63 bei 0% Änderung vom Basiswert jedes vorausgehenden Werts („Precedent“)). Eine positiv geneigte Linie zeigt ein positives Verhältnis oder eine positive Auswirkung an, während eine negativ geneigte Linie ein negatives Verhältnis anzeigt (z.B., Investitionen ist negativ geneigt, was bedeutet, dass je höher das Investitionsniveau, umso geringer der NPV). Der absolute Wert der Neigung zeigt die Größenordnung der Auswirkung (eine steile Linie deutet auf eine höhere Auswirkung auf die NPV y-Achse, gegeben eine Änderung in der vorausgehenden x-Achse).
- ② Das Tornadodiagramm stellt dies auf eine andere graphische Weise dar, wobei die vorausgehende Variable („Precedent“) mit der höchsten Auswirkung als erste aufgelistet ist. Die x-Achse ist der NPV-Wert und das Zentrum des Diagramms ist die Basisfallsituation. Die grünen Balken im Diagramm deuten auf positive Auswirkung, während die roten Balken eine negative Auswirkung anzeigen. Infolgedessen, mit Bezug auf die Investitionen, deuten die roten Balken auf der rechten Seite auf eine negative Auswirkung der Investition auf einem höheren NPV hin—in anderen Worten, Kapitalinvestition und NPV sind negativ korreliert. Das Gegenteil trifft für Preis und Produktmenge A bis C zu (dessen grüne Balken sind auf der rechten Seite des Diagramms).



Bild 5.3—Bericht der Tornadoanalyse

Notes:

Vergessen Sie nicht, dass die Tornadoanalyse eine *statische* Sensibilitätsanalyse ist, die auf jede Inputvariable im Modell angewendet ist—das heißt, jede Variable wird individuell gestört und die resultierenden Auswirkungen werden tabelliert. Dies macht die Tornadoanalyse zu einer Schlüsselkomponente, die vor der Ausführung einer Simulation durchgeführt werden soll. Einer der ersten Schritte in der Risikoanalyse ist der, in dem die wichtigsten Auswirkungstreiber im Modell erfasst und identifiziert werden. Der nächste Schritt ist es zu identifizieren, welche dieser wichtigen Auswirkungstreiber ungewiss sind. Diese ungewissen Auswirkungstreiber sind die kritischen Erfolgstreiber eines Projekts, wobei die Ergebnisse des Modells von diesen kritischen

Erfolgstreibern abhängen. Diese Variablen sind diejenigen, die man simulieren sollte. Verschenden Sie keine Zeit in der Simulation von Variablen, die weder Ungewiss sind noch eine geringe Auswirkung auf die Ergebnisse haben. Tornadodiagramme helfen bei der schnellen und einfachen Identifizierung dieser kritischen Erfolgstreiber. Diesem Beispiel folgend, könnte es sein, dass der Preis und die Menge simuliert sein sollten, angenommen, dass sowohl die erforderliche Investition als auch der effektive Steuersatz im Voraus bekannt und gleichbleibend sind.

Vorausgehende Zelle	Grundwert: 96.6261638553219			Inputänderungen		
	Outputnachteil	Outputvorteil	Effektiver Bereich	Inputnachteil	Inputvorteil	Grundfallwert
C36: Investitionen	276.6261639	-83.373836	360.00	\$1,620.00	\$1,980.00	\$1,800.00
C9: Effektiver Steuersatz	219.7269269	-26.474599	246.20	36.00%	44.00%	40.00%
C12: Prod A Durchschnittspreis	3.425542398	189.82679	186.40	\$9.00	\$11.00	\$10.00
C13: Prod B Durchschnittspreis	16.70663096	176.5457	159.84	\$11.03	\$13.48	\$12.25
C15: Prod A Menge	23.1774976	170.07483	146.90	45.00	55.00	50.00
C16: Prod B Menge	30.5329996	162.71933	132.19	31.50	38.50	35.00
C14: Prod C Durchschnittspreis	40.14658725	153.10574	112.96	\$13.64	\$16.67	\$15.15
C17: Prod C Menge	48.04736933	145.20496	97.16	18.00	22.00	20.00
C5: Markt Risikokorrigierter Diskontsatz	138.2391267	57.029841	81.21	13.50%	16.50%	15.00%
C8: Preiserosionsrate	116.803809	76.640952	40.16	4.50%	5.50%	5.00%
C7: Annualisierte Verkaufswachstumsrate	90.58835433	102.68541	12.10	1.80%	2.20%	2.00%
C24: Abschreibung	95.08417251	98.168155	3.08	\$9.00	\$11.00	\$10.00
C25: Amortisierung	96.16356645	97.088761	0.93	\$2.70	\$3.30	\$3.00
C27: Zinsausgaben	97.08876126	96.163566	0.93	\$1.80	\$2.20	\$2.00

Bild 5.4—Sensibilitätstabelle

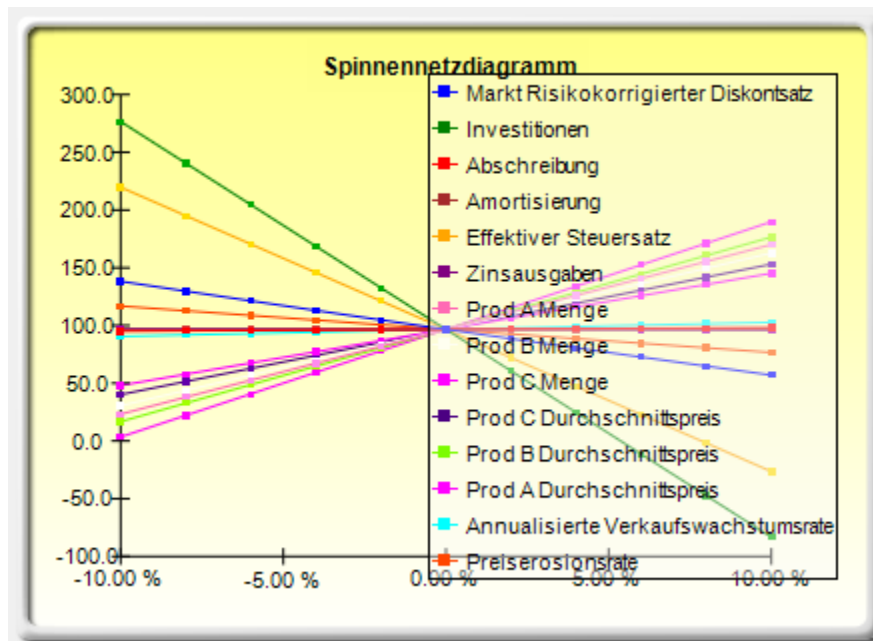


Bild 5.5—Spinnennetzdiagramm

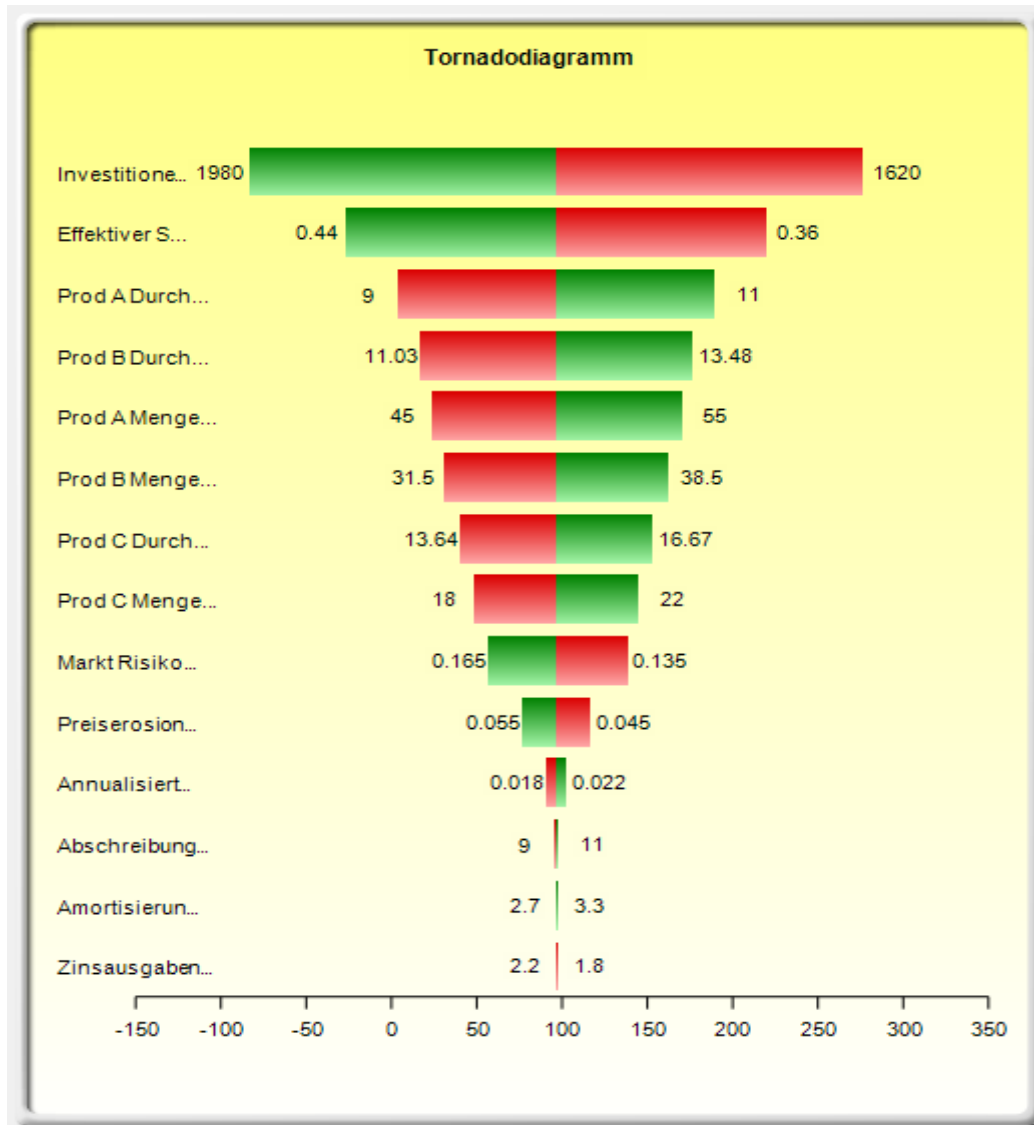


Bild 5.6—Tornadodiagramm

Obwohl das Tornadodiagramm einfacher zu lesen ist, ist das Spinnennetzdiagramm wichtig, um festzustellen, ob irgendwelche Nichtlinearitäten im Modell vorhanden sind. Zum Beispiel, Bild 5.7 zeigt ein weiteres Spinnennetzdiagramm, wo die Nichtlinearitäten ziemlich offensichtlich sind (die Linien auf dem Diagramm sind nicht gerade sondern kurvenförmig). Das verwendete Modell ist *Tornado- und Spinnennetzdiagramme (nicht linear)*, welches das Black-Scholes Optionsbewertungsmodell als Beispiel verwendet. Man kann solche Nichtlinearitäten nicht aus einem Tornadodiagramm feststellen und sie könnten eine wichtige Information im Modell sein oder den Entscheidungsträgern wichtige Erkenntnisse über die Modelldynamik liefern.

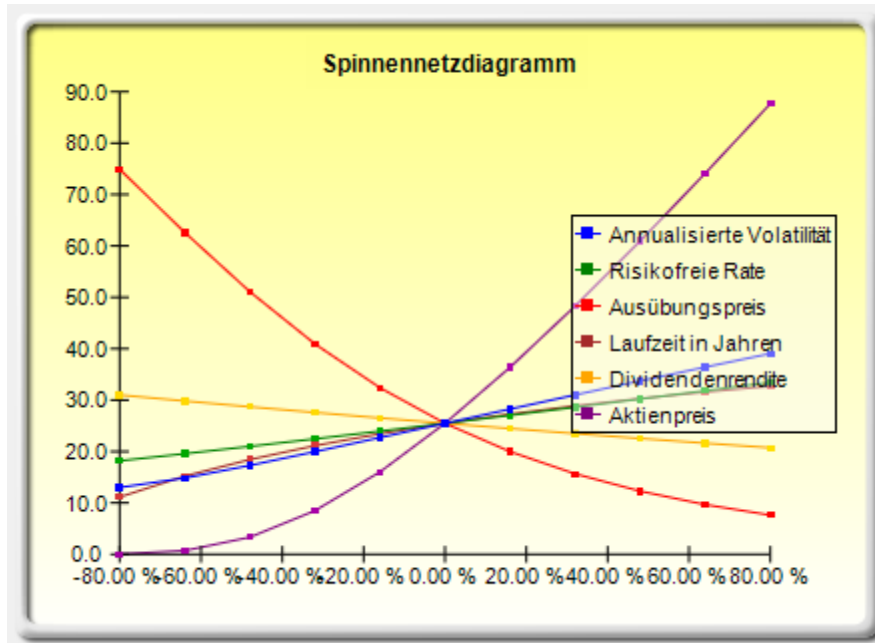


Bild 5.7—Nicht lineares Spinnennetzdiagramm

Zusätzliche Bemerkungen über Tornado:

Bild 5.2 zeigt die Benutzeroberfläche des Tornadoanalysetools. Bitte beachten Sie, dass es, beginnend von der Version 4 von Risiko Simulator, einige neue Verbesserungen gibt. Es folgen einige Tipps über die Ausführung der Tornadoanalyse und weitere Details über diese neuen Verbesserungen:

- Man sollte die Tornadoanalyse nie nur einmal ausführen. Sie dient als ein Modelldiagnosetool, was bedeutet, dass sie idealerweise mehrere Male auf dem gleichen Modell ausgeführt werden sollte. Zum Beispiel, bei einem großen Modell kann man die Tornadoanalyse das erste Mal unter Verwendung aller Standardeinstellungen ausführen und alle vorausgehenden Variablen („Precedents“) sollten angezeigt werden (wählen Sie **Alle Variablen anzeigen**). Daraus könnten sich ein großer Bericht und lange (und potentiell unansehnliche) Tornadodiagramme ergeben. Trotz alledem bietet dies einen hervorragenden Anfangspunkt an, um festzustellen, wie viele der vorausgehenden Variablen („Precedents“) als kritische Erfolgsfaktoren zu betrachten sind (z.B., das Tornadodiagramm könnte anzeigen, dass die ersten 5 Variablen eine höhere Auswirkung auf den Output haben, während die restlichen 200 Variablen keine oder eine geringe Auswirkung ausüben). In diesem Fall wird eine zweite Tornadoanalyse ausgeführt, die weniger Variable anzeigt (z.B., wählen Sie **Die oberen 10 Variablen anzeigen**, wenn die ersten fünf kritisch sind, was einen schönen Bericht und ein Tornadodiagramm erstellt, das einen Kontrast zwischen den Hauptfaktoren und den weniger kritischen Faktoren zeigt. Das heißt, Sie sollten nie ein Tornadodiagramm nur mit den Hauptfaktoren zeigen, ohne auch einige weniger kritische Variablen als einen Kontrast zu ihren Auswirkungen auf den Output anzuzeigen). Als letztes, man kann die Standardtestpunkte von $\pm 10\%$ auf

einen größeren Wert erhöhen, um für Nichtlinearitäten zu testen (das Spinnennetzdiagramm wird die nicht linearen Linien zeigen und die Tornadodiagramme werden zu einer Seite verzerrt sein, wenn die vorausgehenden Auswirkungen nicht linear sind).

- **Zellenadresse verwenden** ist immer eine gute Idee, wenn Ihr Modell groß ist, Diese Option erlaubt es Ihnen, den Speicherort (Arbeitsblattname und Zellenadresse) einer vorausgehenden Zelle zu identifizieren. Wenn diese Option nicht aktiviert ist, wird die Software ihre eigene „Fuzzy-Logik“ anwenden, in einem Versuch den Namen einer vorausgehenden festzustellen (gelegentlich könnten die Namen in einem großen Modell mit wiederholten Variablen unübersichtlich werden oder zu lang sein, was das Tornadodiagramm unansehnlich machen könnte).
- Die Optionen **Dieses Arbeitsblatt analysieren** und **Alle Arbeitsblätter analysieren** erlauben Ihnen nachzuprüfen, ob die vorausgehenden Variablen („Precedents“) nur Teil des aktuellen Arbeitsblatts sein sollten oder alle Arbeitsblätter im gleichen Arbeitsbuch einschließen. Diese Option ist nützlich, wenn Sie nur versuchen einen Output basierend auf Werten im aktuellen Blatt zu analysieren, im Gegenteil zur Ausführung einer globalen Suche aller verknüpften vorausgehenden Variablen („Precedents“) über mehrfache Arbeitsblätter im gleichen Arbeitsbuch.
- **Globale Einstellungen verwenden** ist nützlich, wenn Sie ein großes Modell haben und möchten alle vorausgehenden Variablen („Precedents“) bei, sagen wir, $\pm 50\%$ statt der Standardeinstellung von 10% testen. Statt jeden einzelnen Testwert der vorausgehenden Variablen („Precedent“) einen nach dem anderen ändern zu müssen, können Sie diese Option wählen, eine Einstellung ändern, dann irgendwo anders in der Benutzeroberfläche klicken und die gesamte Liste der vorausgehenden Variablen („Precedents“) wird sich ändern. Die Deaktivierung dieser Option erlaubt Ihnen die Kontrolle über die Änderung der Testpunkte, eine vorausgehende Variable („Precedent“) nach der anderen.
- **Null oder Leerwerte ignorieren** ist eine standardmäßig aktivierte Option, wobei die vorausgehenden Zellen mit Null oder Leerwerten nicht in der Tornadoanalyse ausgeführt werden. Dies ist die typische Einstellung.
- **Höchstmögliche Ganzzahlwerte** ist eine Option, die schnell alle möglichen vorausgehenden Zellen identifiziert, die derzeit Ganzzahlinputs haben. Dies ist manchmal wichtig, wenn Ihr Modell Schalter verwendet (z.B., Funktionen wie WENN (IF) eine Zelle 1 ist, dann geschieht etwas, und WENN (IF) eine Zelle einen Wert von 0 hat, dann geschieht etwas anderes; oder auch Ganzzahlen wie 1, 2, 3 und so weiter, welche Sie nicht testen wollen). Zum Beispiel, $\pm 10\%$ eines Flagschalterwerts von 1 ergibt einen Testwert von 0.9 und 1.1, von denen beide irrelevante und falsche Inputwerte im Modell sind, sodass Excel die Funktion als einen Fehler interpretieren könnte. Diese Option, wenn ausgewählt, wird potentielle Problembereiche für die Tornadoanalyse schnell hervorheben und Sie können feststellen, welche vorausgehenden Variablen („Precedents“) Sie manuell ein- oder ausschalten müssen. Sie können aber auch die

Option **Alle möglichen Ganzzahlwerte ignorieren**, um alle diese gleichzeitig auszuschalten.

Sensibilitätsanalyse

Theorie:

Eine verwandte Funktion ist die Sensibilitätsanalyse. Während die Tornadoanalyse (Tornadodiagramme und Spinnennetzdiagramme) statische Störungen *vor* einem Simulationslauf anwendet, wendet die Sensibilitätsanalyse dynamische Störungen an, die *nach* dem Simulationslauf erstellt wurden. Tornado- und Spinnennetzdiagramme sind die Ergebnisse von statischen Störungen. Das heißt, jede vorausgehende Variable („Precedent“) oder Hypothesenvariable wird eine nach der anderen um eine vordefinierte Menge gestört und die Fluktuationen in den Ergebnissen werden tabelliert. Sensibilitätsdiagramme dagegen sind die Ergebnisse von dynamischen Störungen, im Sinne, dass mehrfache Hypothesen gleichzeitig gestört und ihre Interaktionen im Modell und ihre Korrelationen zwischen Variablen in den Fluktuationen der Ergebnisse erfasst werden. Tornadodiagramme identifizieren deshalb die Variablen, welche die Ergebnisse am meisten treiben und welche daher zur Simulation geeignet sind. Sensibilitätsdiagramme dagegen identifizieren die Auswirkung auf die Ergebnisse, wenn mehrfachen interagierenden Variablen zusammen im Modell simuliert werden. Diese Auswirkung ist deutlich im Bild 5.8 dargestellt. Bitte bemerken Sie, dass die Rangliste der kritischen Erfolgstreiber ähnlich wie bei dem Tornadodiagramm in den vorherigen Beispielen ist. Wenn man allerdings Korrelationen zwischen den Hypothesen hinzufügt, zeigt das Bild 5.9 eine sehr unterschiedliche Situation. Bitte bemerken Sie zum Beispiel, dass die Preiserosion eine geringe Auswirkung auf NPV hatte. Wenn aber einige der Inputhypothesen korreliert werden, wird die Preiserosion, auf Grund der Interaktion, die zwischen diesen korrelierten Variablen existiert, eine größere Auswirkung haben.

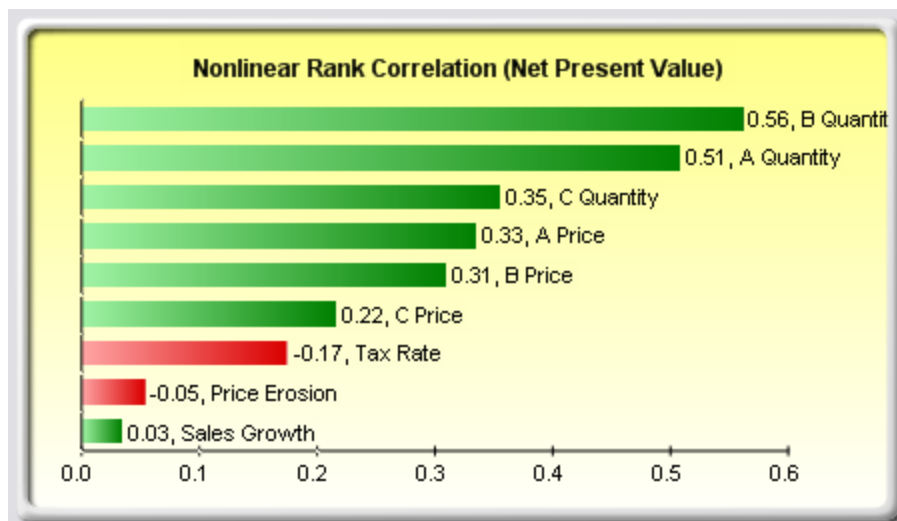


Bild 5.8—Sensibilitätsdiagramm ohne Korrelationen

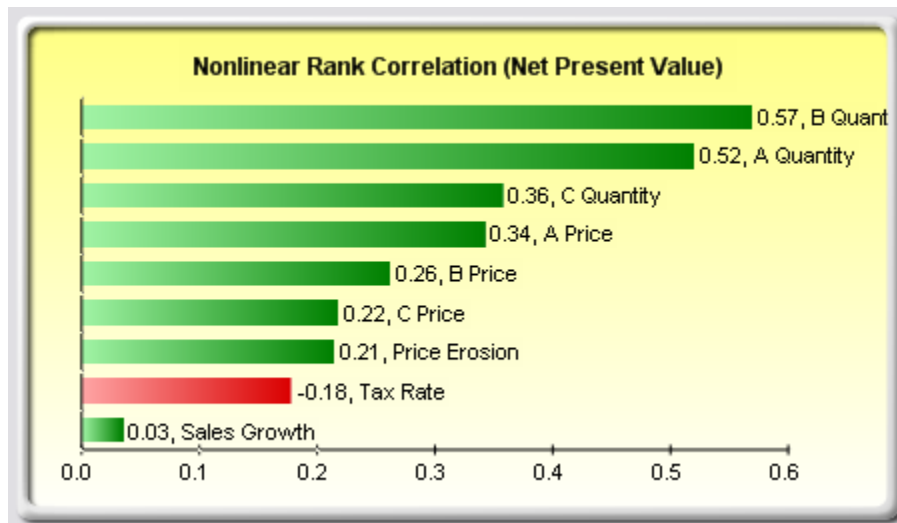


Bild 5.9—Sensibilitätsdiagramm mit Korrelationen

Prozedur:

1. Öffnen Sie oder kreieren sie ein Modell, definieren Sie Hypothesen und Vorausberechnungen und führen Sie die Simulation aus (dieses Beispiel verwendet die Datei *Tornado- und Sensibilitätsdiagramme (lineare)*)
2. Wählen Sie **Risiko Simulator | Tools | Sensibilitätsanalyse**
3. Wählen Sie die gewünschte zu analysierende Vorausberechnung aus und klicken Sie auf **OK** (Bild 5.10)

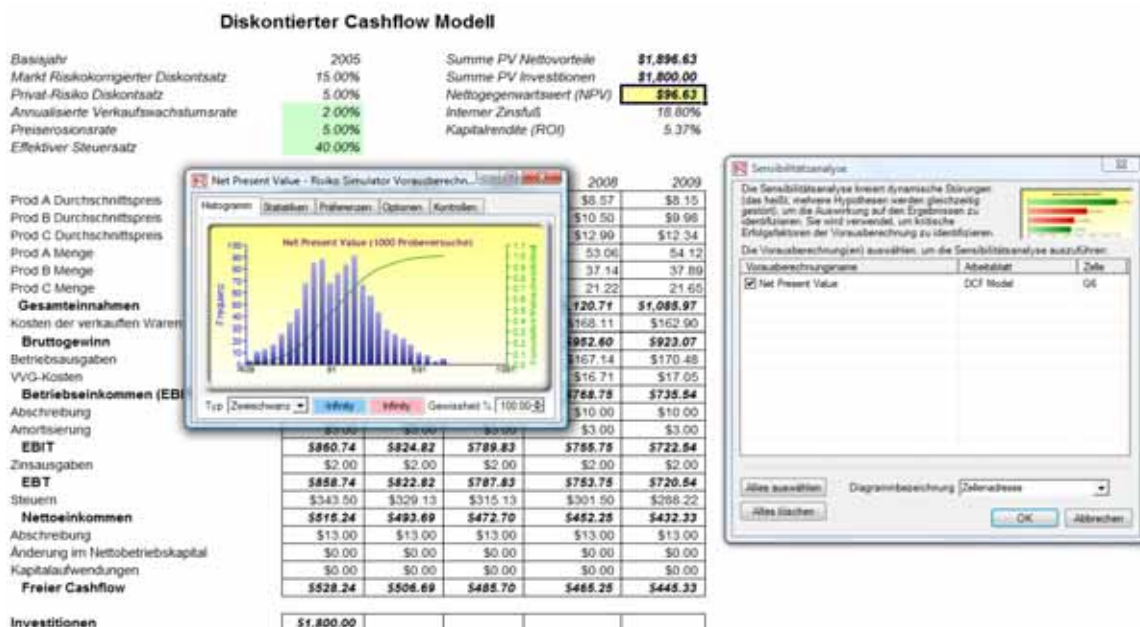


Bild 5.10—Eine Sensibilitätsanalyse ausführen

Interpretierung der Ergebnisse:

Die Ergebnisse der Sensibilitätsanalyse beinhalten einen Bericht und zwei Schlüsseldiagramme. Das erste ist ein Diagramm der nicht linearen Rangkorrelation (Bild 5.11), welches die Korrelationspaare der Hypothesen-Vorausberechnung vom höchsten zum niedrigsten auflistet. Diese Korrelationen sind nicht linear and nicht parametrisch, was sie von allen Verteilungsanforderungen befreit (das heißt, man kann eine Hypothese mit einer Weibull-Verteilung mit einer anderen Hypothese mit einer Betaverteilung vergleichen). Die Ergebnisse dieses Diagramms sind ziemlich ähnlich denen der vorher beschriebenen Tornadoanalyse (natürlich ohne den Kapitalinvestitionswert, den wir als bekannten Wert bezeichnet hatten und der deshalb nicht simuliert wurde), mit einer speziellen Ausnahme. Der Steuersatz wurde zu einer viel niedrigeren Position im Sensibilitätsanalysediagramm (Bild 5.11) im Vergleich zum Tornadodiagramm (Bild 5.6) heruntergestuft. Das liegt daran, dass der Steuersatz an sich eine signifikante Auswirkung hat, aber sobald die anderen Variablen im Modell interagieren, scheint es als ob der Steuersatz eine weniger dominante Auswirkung hat (das liegt daran, dass der Steuersatz eine kleinere Verteilung hat, da historische Steuersätze nicht dazu neigen, zu viel zu schwanken und auch weil der Steuersatz ein direkter Prozentwert der Einnahmen vor Steuern ist, worauf andere vorausgehende Variablen eine größere Auswirkung haben). Dieses Beispiel beweist, dass die Ausführung einer Sensibilitätsanalyse nach einem Simulationslauf wichtig ist, um festzustellen, ob es irgendwelche Interaktionen im Modell gibt und ob die Auswirkungen auf bestimmte Variablen immer noch halten. Das zweite Diagramm (Bild 5.12) erläutert die erklärte Prozentvariation. Das heißt, von den Fluktuationen in der Vorausberechnung, wie viel von der Variation kann man durch jede der Hypothesen, nach Berücksichtigung aller Interaktionen zwischen den Variablen, erklären? Bitte bemerken Sie, dass die Summe aller erklärten Variationen normalerweise nah bei 100% liegt (gelegentlich gibt es andere Elemente, die Auswirkungen auf das Modell haben, die aber hier nicht direkt erfasst werden können) und wenn Korrelationen vorhanden sind, könnte die Summe manchmal 100% überschreiten (auf Grund der Interaktionsauswirkungen, die kumulativ sind).

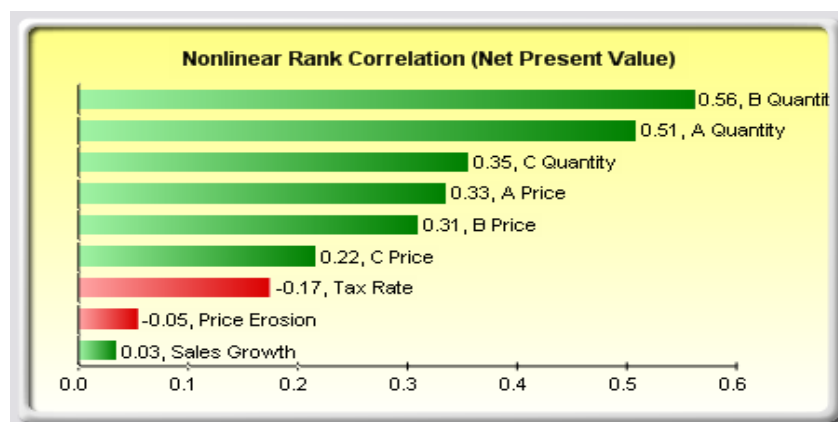


Bild 5.11—Rangkorrelationsdiagramm

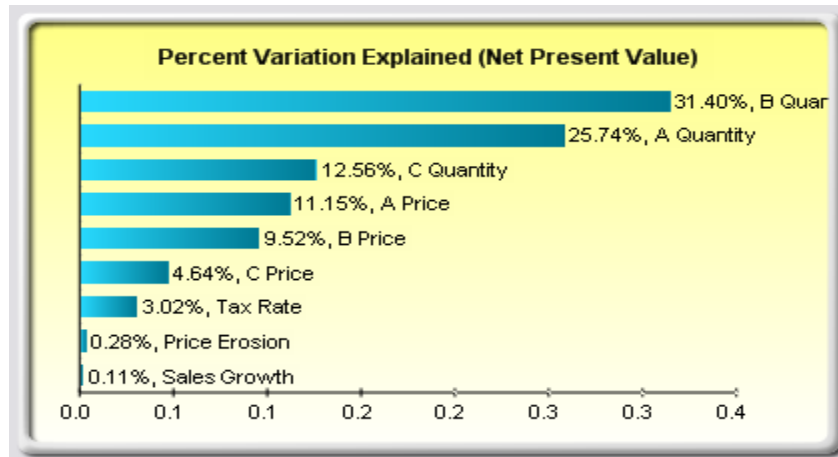


Bild 5.12—Beitrag zum Varianzdiagramm

Bemerkungen:

Die Tornadoanalyse wird vor einem Simulationslauf ausgeführt, während die Sensibilitätsanalyse nach einem Simulationslauf durchgeführt wird. Spinnennetzdiagramme in der Tornadoanalyse können Nichtlinearitäten berücksichtigen, während Rangkorrelationsdiagramme in der Sensibilitätsanalyse nichtlineare und verteilungsfreie Bedingungen nachweisen können.

Verteilungsanpassung: Einzel-Variable und Mehrfach-Variablen

Theorie:

Ein anderes leistungsstarke Simulationstool ist die Verteilungsanpassung. Das heißt, welche Verteilung soll ein Analyst für eine bestimmte Inputvariable in einem Modell verwenden? Welche sind die relevanten Verteilungsparameter? Wenn keine historischen Daten vorhanden sind, dann muss der Analyst Hypothesen über die fraglichen Variablen aufbauen. Eine Behandlungsweise ist die Verwendung der Delphi-Methode, wo eine Expertengruppe mit der Schätzung des Verhaltens jeder Variablen beauftragt ist. Zum Beispiel, man kann eine Gruppe von Maschinenbauingenieuren mit der Auswertung der extremen Möglichkeiten des Diameters einer Schraubenfeder mittels rigoroses Experimentierens oder Daumenschätzungen beauftragen. Diese Werte können als die Inputparameter der Variablen verwendet werden (z.B., eine Uniformverteilung mit extremen Werten zwischen 0.5 und 1.2). Wenn das Testen nicht möglich ist (z.B., Marktanteil und Einnahmenwachstumsrate), kann das Management dennoch Schätzungen über die potentiellen Ausgänge machen und die Szenarien für den besten, den wahrscheinlichsten und den schlimmsten Fall zur Verfügung stellen.

Wenn allerdings zuverlässige historische Daten vorhanden sind, kann man die Verteilungsanpassung benutzen. Angenommen, dass die historische Muster weiter halten und dass die Geschichte dazu neigt, sich zu wiederholen, kann man historische Daten dazu verwenden, um die bestpassende Verteilung mit ihren relevanten Parametern zu finden, um die zu

simulierenden Variablen besser zu definieren. Bilder 5.13 bis 5.15 erläutern das Beispiel einer Verteilungsanpassung. Diese Erklärung verwendet die Datei **Datenanpassung** im Ordner Beispiele.

Prozedur:

- Öffnen Sie ein Tabellenblatt mit existierenden Daten zur Anpassung
- Wählen Sie die Daten aus, die Sie anpassen möchten (die Daten sollten in einer einzelnen Spalte mit mehreren Reihen sein)
- Wählen Sie **Risiko Simulator | Tools | Verteilungsanpassung (Einzel-Variable)**
- Wählen Sie die spezifischen Verteilungen, auf welche Sie anpassen möchten oder behalten Sie die Standardeinstellung, wobei alle Verteilungen ausgewählt sind und klicken Sie auf **OK** (Bild 5.13)
- Überprüfen Sie die Ergebnisse der Anpassung, wählen Sie die gewünschte relevante Verteilung und klicken Sie auf **OK** (Bild 5.14)

	A	B	C	D	E	F	G	H
1	Normal		Var X	Var Y	Var Z			
2	93.75		87.53	45.29	6.00			
3	109.52		99.66	46.94	6.00			
4	101.17		108.75	45.96	6.00			
5	102.29		87.41	52.09	8.00			
6	105.58							
7	99.55							
8	86.79							
9	105.20							
10	113.63							
11	105.90							
12	90.68							
13	96.20							
14	79.74							
15	91.49							
16	98.28							
17	97.70							
18	97.85							
19	93.73							
20	92.06							
21	85.51							
22	103.21							
23	87.45							
24	96.40							
25	92.41							
26	82.75							
27	103.65							
28	90.19							
29	112.42							
30	103.22							
31	91.56							
32	86.04							
33	115.40							

Einzel-Anpassung

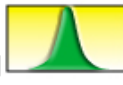
Die Verteilungsanpassung nimmt existierende Rohdaten und findet statistisch die best passende Verteilung (nämlich durch die Optimierung der Parameter jeder Verteilung und die Ausführung von statistischen Hypothesentests).

Verteilungstyp

☒ An kontinuierliche Verteilungen anpassen
 ☐ An diskrete Verteilungen anpassen

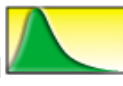
Verteilungen zum Anpassen auswählen:

☒ 
Beta

☒ 
Cauchy

☒ 
Chi-Quadrat

☒ 
Exponentielle

☒ 
F

☒ 
Gamma

			95.27	53.87	6.00
			95.33	46.24	7.00
			102.26	54.89	6.00

Bild 5.13—Einzel-Variable Verteilungsanpassung

Interpretierung der Ergebnisse:

Die zum Testen bestimmte Nullhypothese ist so, dass die angepasste Verteilung dieselbe Verteilung ist, wie die Bevölkerung von der die anzupassenden Stichprobendaten stammen. Deshalb, wenn der berechnete p-Wert kleiner als ein kritisches Alphaniveau (typischerweise 0.10 oder 0.05) ist, dann ist die Verteilung die falsche Verteilung. Indessen, je höher der p-Wert, umso besser passt sich die Verteilung den Daten an. Grob gesagt können Sie den p-Wert als einen *erklärten Prozentsatz* betrachten. Das heißt, wenn der p-Wert 0.9727 ist (Bild 5.14), dann wird die Einstellung einer Normalverteilung mit einem Mittelwert von 99.28 und einer Standardabweichung von 10.17 zirka 97.27% der Variation in den Daten erklären, was auf eine besonders gute Anpassung hindeutet. Sowohl die Ergebnisse (Bild 5.14) als auch der Bericht (Bild 5.15) zeigen die Teststatistiken, den p-Wert, die theoretischen Statistiken (basierend auf der ausgewählten Verteilung), die empirischen Statistiken (basierend auf den Rohdaten), die Originaldaten (um einen Nachweis der verwendeten Daten zu bewahren) und die Hypothesen, komplett mit den relevanten Verteilungsparametern (das heißt, ob sie die Option zur automatische Generierung der Hypothesen ausgewählt haben und ob ein Simulationsprofil schon vorhanden ist). Die Ergebnisse stellen auch eine Rangliste aller ausgewählten Verteilungen auf und zeigen wie gut sie sich den Daten anpassen.

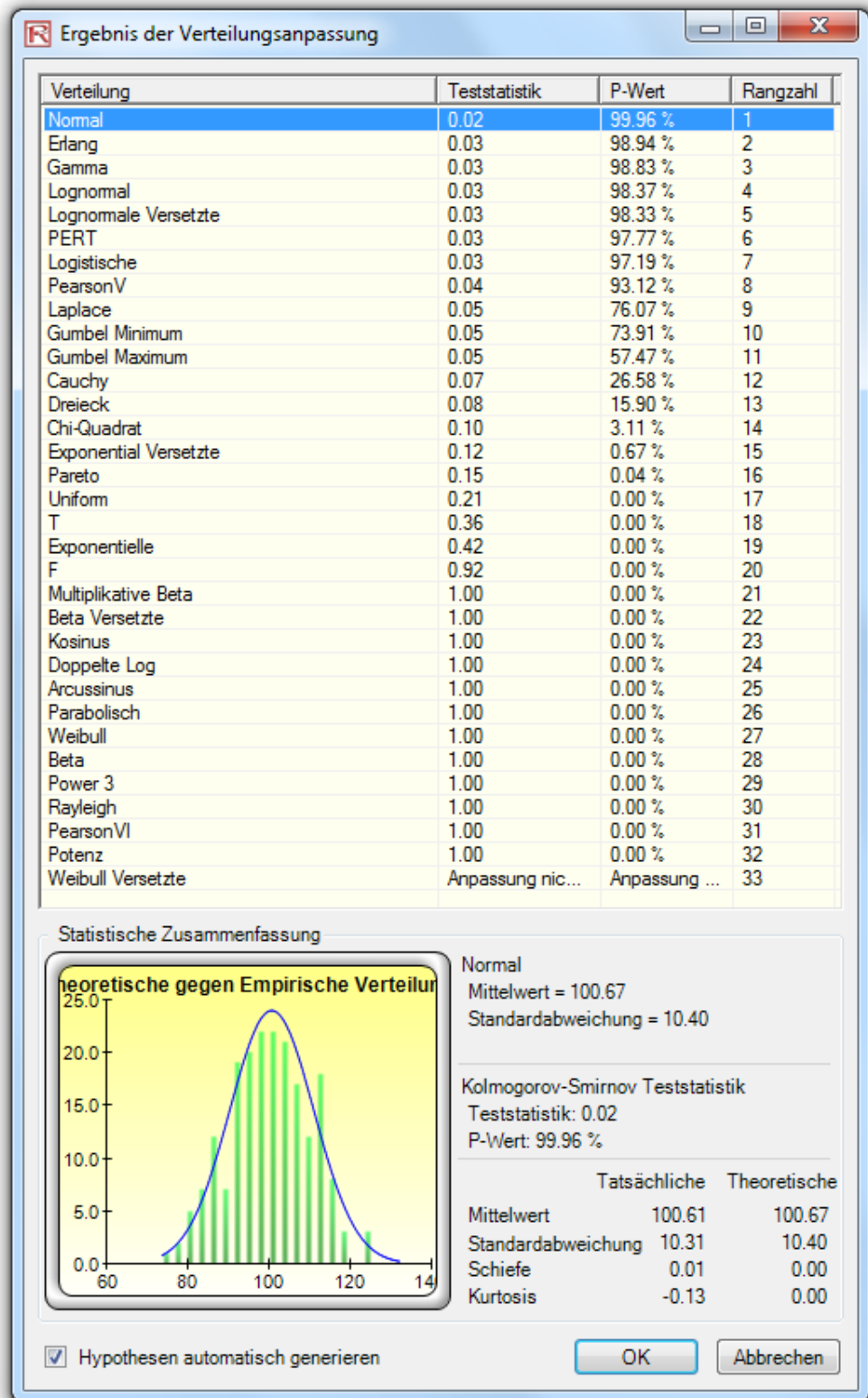


Bild 5.14 - Ergebnisse der Verteilungsanpassung

Einzelvariableverteilungsanpassung

Statistische Zusammenfassung

Angepasste Hypothese **100.61**

Angepasste Verteilung **Normal**

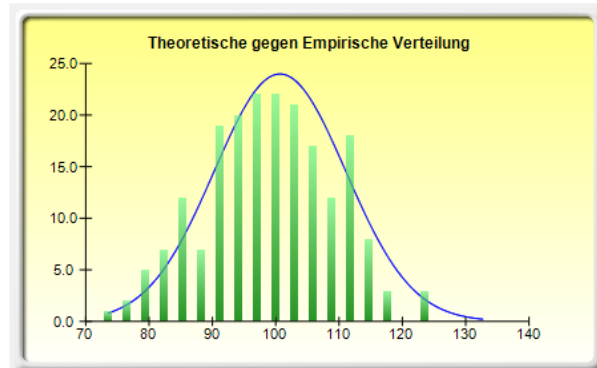
Mittelwert 100.67

Sigma 10.40

Kolmogorov-Smirnov Statistik 0.02

P-Wert für Teststatistik 0.9996

	Tatsächliche	Theoretische
Mittelwert	100.61	100.67
Standardabweichung	10.31	10.40
Schiefe	0.01	0.00
Exzess Kurtosis	-0.13	0.00



Originale angepasste Daten

73.53	78.21	78.52	79.50	79.72	79.74	81.56	82.08	82.68	82.75	83.34	83.64	84.09
84.66	85.00	85.35	85.51	86.04	86.79	86.82	86.91	87.02	87.03	87.45	87.53	87.66
88.05	88.45	88.51	89.95	90.19	90.54	90.68	90.96	91.25	91.49	91.56	91.94	92.06
92.36	92.41	92.45	92.70	92.80	92.84	93.21	93.26	93.48	93.73	93.75	93.77	93.82
94.00	94.15	94.51	94.57	94.64	94.69	94.95	95.57	95.62	95.71	95.78	95.83	95.97
96.20	96.24	96.40	96.43	96.47	96.81	96.88	97.00	97.07	97.21	97.23	97.48	97.70
97.77	97.85	98.15	98.17	98.24	98.28	98.32	98.33	98.35	98.65	99.03	99.27	99.46
99.47	99.55	99.73	99.96	100.08	100.24	100.36	100.42	100.44	100.48	100.49	100.83	101.17
101.28	101.34	101.45	101.46	101.55	101.73	101.74	101.81	102.29	102.55	102.58	102.60	102.70
103.17	103.21	103.22	103.32	103.34	103.45	103.65	103.66	103.72	103.81	103.90	103.99	104.46
104.57	104.76	105.20	105.44	105.50	105.52	105.58	105.66	105.87	105.90	105.90	106.29	106.35
106.59	107.01	107.68	107.70	107.93	108.17	108.20	108.34	108.42	108.43	108.49	108.70	109.15
109.22	109.35	109.52	109.75	110.04	110.16	110.25	110.54	111.05	111.06	111.44	111.76	111.90
111.95	112.07	112.19	112.29	112.32	112.42	112.48	112.85	112.92	113.50	113.59	113.63	113.70
114.13	114.14	114.21	114.91	114.95	115.40	115.58	115.66	116.58	116.98	117.60	118.67	119.24
119.52	124.14	124.16	124.39	132.30								

Bild 5.15 - Bericht der Verteilungsanpassung

Um mehrfache Variablen anzupassen, ist der Prozess ziemlich ähnlich dem der Anpassung von individuellen Variablen. Allerdings sollten die Daten in Spalten geordnet sein (das heißt, jede Variable ist als Spalte geordnet) und alle Variablen werden eine nach der anderen angepasst.

Prozedur:

- 📄 Ein Tabellenblatt mit existierenden Daten zur Anpassung öffnen
- 📄 Wählen Sie die gewünschten anzupassenden Daten (die Daten sollten in mehrfachen Spalten mit mehrfachen Reihen sein)
- 📄 Wählen Sie **Risiko Simulator | Tools | Verteilungsanpassung (Mehrfach-Variablen)**
- 📄 Überprüfen Sie die Daten, wählen Sie die gewünschten relevanten Verteilungstypen und klicken Sie auf **OK**

Bemerkung:

Bitte bemerken Sie, dass die statistischen Rangreihenmethoden, welche in den Verteilungsanpassungsroutinen verwendet werden, der Chi-Quadrat-Test und der Kolmogorov-Smirnov-Test sind. Der erste wird verwendet, um diskrete Verteilungen und der zweite, um kontinuierliche Verteilungen zu testen. Kurz gefasst, ein Hypothesentest test verbunden mit einer internen Optimierungsroutine wird verwendet, um die bestpassenden Parameter für jeder der

getesteten Verteilungen zu finden und die Ergebnisse werden von der besten zur schlechtesten Anpassung geordnet

Bootstrap-Simulation

Theorie:

Die Bootstrap-Simulation ist eine einfache Methode, welche die Zuverlässigkeit oder Genauigkeit der Vorausberechnungsstatistiken oder anderer Stichprobenrohdaten schätzt. Im Wesentlichen wird die Bootstrap-Simulation in Hypothesentesten verwendet. Klassische in der Vergangenheit verwendete Methoden stützten sich auf mathematischen Formeln, um die Genauigkeit der Stichprobenstatistiken zu beschreiben. Diese Methoden nehmen an, dass die Verteilungen einer Stichprobenstatistik sich einer Normalverteilung naht, was die Berechnung des Standardfehlers oder des Konfidenzintervalls der Statistik ziemlich einfach macht. Wenn jedoch die Stichprobenverteilung einer Statistik nicht normalverteilt ist oder leicht gefunden werden kann, sind diese klassischen Methoden schwer verwendbar oder ungültig. Bootstrapping dagegen analysiert Stichprobenstatistiken auf empirische Weise: Die Daten werden einem Stichprobenverfahren wiederholt unterworfen und es werden Verteilungen der verschiedenen Statistiken aus jedem Stichprobenverfahren erstellt.

Prozedur:

- Führen sie eine Simulation aus
- Wählen Sie **Risiko Simulator | Tools | Nicht parametrischer Bootstrap**
- Wählen Sie nur eine Vorausberechnung für das Bootstrap-Verfahren, wählen Sie die Statistik(en) für das Bootstrap-Verfahren, geben Sie die Anzahl der Bootstrap-Probeversuche ein und klicken sie auf **OK** (Bild 5.16)

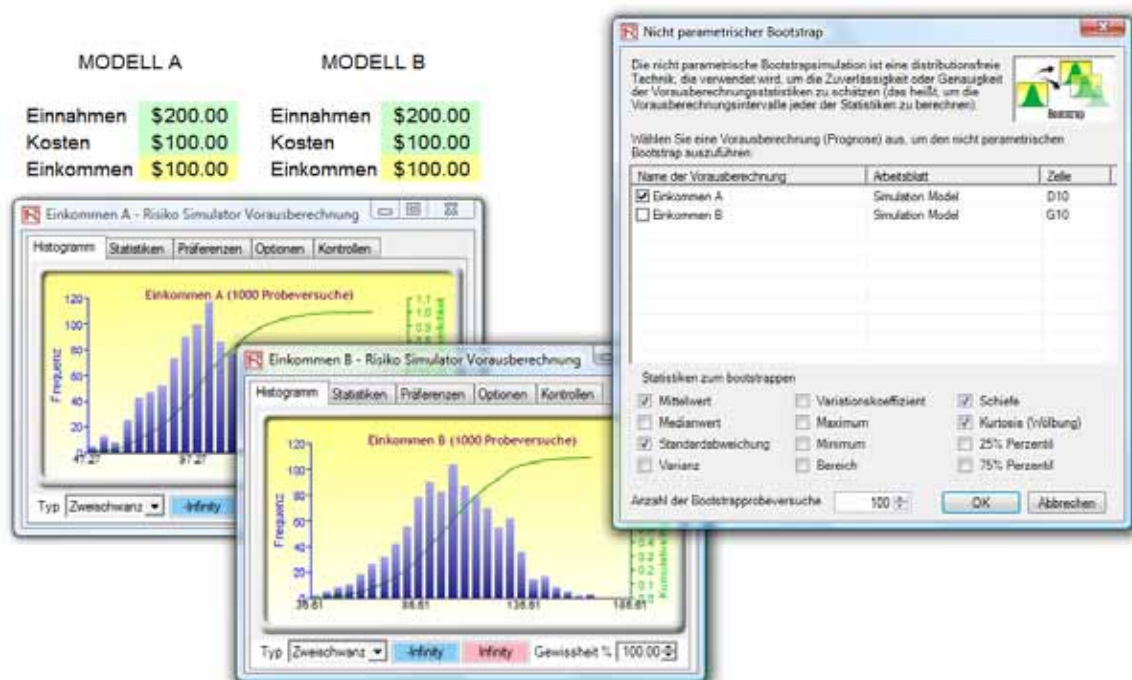


Bild 5.16—Nicht parametrische Bootstrap-Simulation

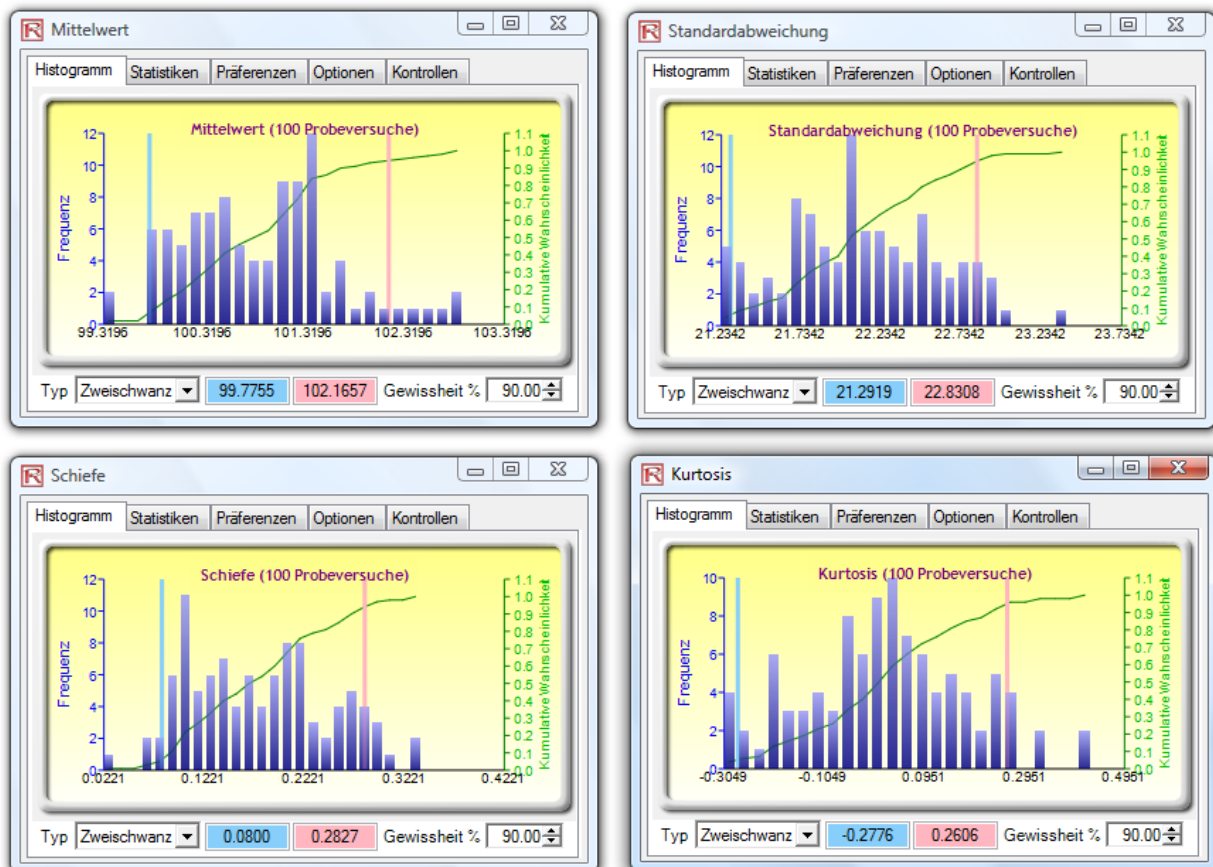


Bild 5.17—Ergebnisse der Bootstrap-Simulation

Interpretierung der Ergebnisse:

Im Wesentlichen kann man die nicht parametrische Bootstrap-Simulation als *eine Simulation basierend auf einer Simulation* bezeichnen. Folglich, nach Ausführung einer Simulation werden die resultierenden Statistiken visualisiert, aber die Genauigkeit dieser Statistiken und ihre statistische Signifikanz sind gelegentlich bedenklich. Zum Beispiel, wenn die Schiefest Statistik eines Simulationslaufs -0.10 ist, ist diese Verteilungen wirklich negativ verzerrt oder der leicht negative Wert ist auf reinen Zufall zurückzuführen? Wie ist es mit -0.15 , -0.20 und so weiter? In anderen Worten, wie weit ist weit genug, sodass man diese Verteilung als negativ verzerrt betrachten kann? Dieselbe Frage ist auf alle anderen Statistiken anwendbar. Ist eine Verteilung statistisch identisch mit einer anderen Verteilung bezüglich irgendwelcher berechneten Statistiken oder sind sie signifikant unterschiedlich? Bild 5.17 stellt einige Beispiele von Ergebnissen aus einem Bootstrap dar. Zum Beispiel, die 90% Konfidenz für die Schiefest Statistik liegt zwischen -0.0189 und 0.0952 , sodass der Wert von 0 innerhalb dieser Konfidenz fällt, was darauf hindeutet, dass bei einer 90% Konfidenz, die Schiefe dieser Vorausberechnung nicht statistisch signifikant abweichend von Null ist, oder dass man diese Verteilung als symmetrisch und nicht verzerrt betrachtet. Umgekehrt, wenn der Wert 0 außerhalb dieser Konfidenz fällt, ist das Gegenteil wahr: Die Verteilung ist verzerrt (positiv verzerrt, wenn die Vorausberechnungsstatistik positiv ist, und negativ verzerrt, wenn die Vorausberechnungsstatistik negativ ist).

Bemerkungen:

Der Begriff *Bootstrap* stammt von der Redensart “to pull oneself up by one’s own bootstraps” (sich an den eigenen Stiefelriemen aus dem Sumpf ziehen) und ist anwendbar, weil diese Methode die Verteilung der Statistiken selber verwendet, um die Genauigkeit der Statistiken zu analysieren. Eine nicht parametrische Simulation ist lediglich das zufällige Herausnehmen mit Zurücklegung von Golfbällen aus einem großen Korb, wobei jeder Golfball auf einem historischen Datenpunkt basiert. Nehmen wir an, dass es 365 Golfbälle im Korb gibt (diese repräsentieren 365 historische Datenpunkte). Stellen Sie sich bitte vor, dass der Wert von jedem zufällig herausgenommenen Golfball auf eine große Weißwandtafel geschrieben wird. Die Ergebnisse der 365 mit Zurücklegung herausgenommenen Bälle sind in der ersten Spalte der Tafel mit 365 Reihen von Zahlen geschrieben. Die relevanten Statistiken (z.B., Mittelwert, Medianwert, Standardabweichung und so weiter) werden auf diese 365 Reihen berechnet. Der Prozess wird dann, sagen wir, 5000 Male wiederholt. Die Weißwandtafel wird jetzt mit 365 Reihen und 5000 Spalten gefüllt sein. Demzufolge werden 5000 Sätze von Statistiken (das heißt, es gibt 5000 Mittelwerte, 5000 Medianwerte, 5000 Standardabweichungen und so weiter) tabelliert und ihre Verteilungen visualisiert. Dann werden die relevanten *Statistiken der Statistiken* tabelliert, wobei man von diesen Ergebnissen feststellen kann, wie „vertrauensvoll“ die simulierten Statistiken sind. In anderen Worten, sagen wir, dass bei einer einfachen Simulation von 10000 Probeversuchen, es sich ergibt, dass der resultierende Vorausberechnungsdurchschnitt bei \$5.00 liegt. Wie sicher ist sich der Analyst über die Ergebnisse? Bootstrapping erlaubt dem Benutzer das Konfidenzintervall der berechneten

Mittelwertstatistik festzustellen, was auf der Verteilung der Statistiken hindeutet. Zum Schluss, Die Ergebnisse des Bootstrap-Verfahrens sind wichtig, weil, in Übereinstimmung mit dem *Gesetz der großen Zahlen* und dem *Theorem des zentralen Grenzwertsatzes* der Statistiktheorie, der Mittelwert der Stichprobenmittelwerte ein unverzerrter Schätzer ist und dem wahren Bevölkerungsmittelwert naht, wenn die Stichprobengröße wächst.

Hypothesentest

Theorie:

Ein Hypothesentest wird durchgeführt, wenn man die Mittelwerte und Varianzen von zwei Verteilungen testet, um festzustellen, ob sie statistisch identisch oder statistisch abweichend voneinander sind. Das bedeutet, um zu sehen, ob die sich ereignenden Unterschiede zwischen den Mittelwerten und Varianzen von zwei verschiedenen Vorausberechnungen auf reinem Zufall basieren oder ob sie tatsächlich statistisch signifikant abweichend voneinander sind.

Prozedur:

- ☑ Eine Simulation ausführen
- ☑ Wählen Sie **Risiko Simulator | Tools | Hypothesentesten**
- ☑ Wählen Sie nur *zwei* auf einmal zu testende Vorausberechnungen, wählen Sie den Typ des gewünschten auszuführenden Hypothesentests und klicken Sie auf **OK** (Bild 5.18)

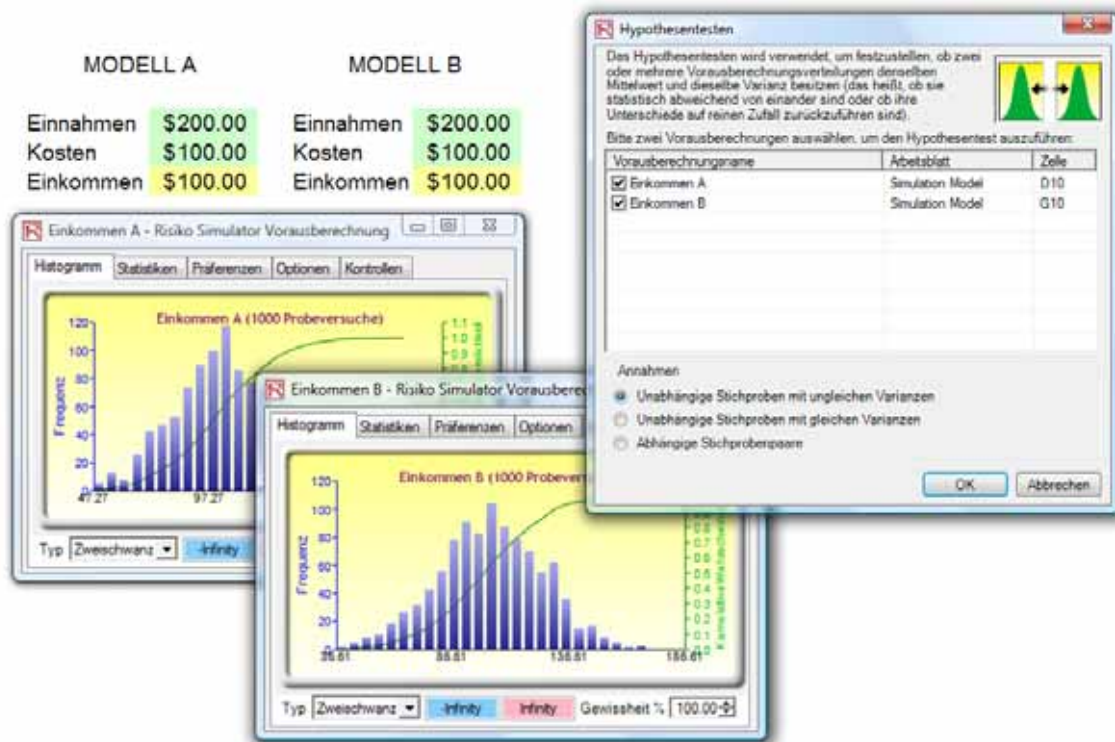


Bild 5.18—Hypothesentest

Interpretierung der Ergebnisse:

Ein Zweiseitiger Hypothesentest wird auf der Nullhypothese (H_0) ausgeführt, sodass die Bevölkerungsmittelwerte der zwei Variablen statistisch identisch miteinander sind. Die Alternativhypothese (H_a) ist, dass die Bevölkerungsmittelwerte statistisch abweichend voneinander sind. Wenn die berechneten p-Werte weniger als oder gleich 0.01, 0.05 oder 0.10 sind, bedeutet es, dass die Nullhypothese abgelehnt wird, was andeutet, dass die Vorausberechnungsmittelwerte statistisch signifikant unterschiedlich an den 1%, 5% und 10% Signifikanzniveaus sind. Wenn die Nullhypothese bei hohen p-Werten nicht abgelehnt wird, sind die Mittelwerte der zwei Vorausberechnungsverteilungen statistisch miteinander ähnlich. Dieselbe Analyse wird einzeln auf die Varianzen von zwei Vorausberechnungen unter Verwendung des Paarweisen F-Tests ausgeführt. Wenn die p-Werte klein sind, dann sind die Varianzen (und Standardabweichungen) statistisch abweichend voneinander; sonst, für große p-Werte, sind die Varianzen statistisch identisch miteinander.

Hypothesentest auf den Mittelwerten und die Varianzen von zwei Vorausberechnungen

Statistische Zusammenfassung

Ein Hypothesentest wird durchgeführt, wenn man die Mittelwerte und Varianzen von zwei Verteilungen testet, um festzustellen, ob sie statistisch identisch oder statistisch abweichend voneinander sind. Das bedeutet, um zu sehen, ob die sich ereignenden Unterschiede zwischen zwei Mittelwerten und zwei Varianzen auf reinem Zufall basieren oder ob sie tatsächlich abweichend voneinander sind. Der Zweivariablen t-Test mit ungleichen Varianzen (man vermutet, dass die Bevölkerungsvarianz der Vorausberechnung 1 abweichend von der Bevölkerungsvarianz der Vorausberechnung 2 ist) ist angemessen, wenn die Vorausberechnungsverteilungen von verschiedenen Bevölkerungen stammen (z.B., Daten, die von zwei verschiedenen geographischen Standorten, zwei verschiedenen Betriebsgeschäftseinheiten, und so weiter, gesammelt wurden). Der Zweivariablen t-Test mit gleichen Varianzen (man vermutet, dass die Bevölkerungsvarianz der Vorausberechnung 1 gleich der Bevölkerungsvarianz der Vorausberechnung 2 ist) ist angemessen, wenn die Vorausberechnungsverteilungen von ähnlichen Bevölkerungen stammen (z.B., Daten, die von zwei verschiedenen Motorenentwürfen mit ähnlichen Spezifikationen, und so weiter, gesammelt wurden). Der gepaarte abhängige Zweivariablen t-Test ist angemessen, wenn die Vorausberechnungsverteilungen von ähnlichen Bevölkerungen stammen (z.B., Daten, die von derselben Gruppe von Kunden aber bei verschiedenen Gelegenheiten, und so weiter, gesammelt wurden).

Ein Zweiseitiger Hypothesentest wird auf der Nullhypothese H_0 ausgeführt, sodass die Bevölkerungsmittelwerte der zwei Variablen statistisch identisch sind. Die Alternativhypothese ist, dass die Bevölkerungsmittelwerte statistisch abweichend voneinander sind. Wenn die berechneten p-Werte weniger als oder gleich 0,01, 0,05 oder 0,10 sind, bedeutet es, dass die Hypothese abgelehnt wird, was andeutet, dass die Vorausberechnungsmittelwerte statistisch erheblich unterschiedlich an den 1%, 5% und 10% Signifikanzniveaus sind. Wenn die Nullhypothese bei hohen p-Werten nicht abgelehnt wird, sind die Mittelwerte der zwei Vorausberechnungsverteilungen statistisch ähnlich. Dieselbe Analyse wird einzeln auf die Varianzen von zwei Vorausberechnungen unter Verwendung des Paarweisen F-Tests ausgeführt. Wenn die p-Werte klein sind, dann sind die Varianzen (und die Standardabweichungen) statistisch abweichend; sonst, für große p-Werte, sind die Varianzen statistisch identisch.

Ergebnis

Hypothesentestannahme:	Ungleiche Varianzen:
Berechnete t-Statistik:	1.015722
P-Wert für die for t-Statistik:	0.309885
Berechnete F-Statistik:	1.063476
P-Wert für die for F-Statistik:	0.330914

Bild 5.19—Ergebnisse des Hypothesentests

Bemerkungen:

Der Zweivariablen t-Test mit ungleichen Varianzen (man erwartet, dass die Bevölkerungsvarianz der Vorausberechnung 1 abweichend von der Bevölkerungsvarianz der Vorausberechnung 2 ist) ist angemessen, wenn die Vorausberechnungsverteilungen von verschiedenen Bevölkerungen

stammen (z.B., Daten, die von zwei verschiedenen Standorten, zwei verschiedenen Betriebsgeschäftseinheiten, und so weiter, gesammelt wurden). Der Zweivariablen t-Test mit gleichen Varianzen (man erwartet, dass die Bevölkerungsvarianz der Vorausberechnung 1 gleich der Bevölkerungsvarianz der Vorausberechnung 2 ist) ist angemessen, wenn die Vorausberechnungsverteilungen von ähnlichen Bevölkerungen stammen (z.B., Daten, die von zwei verschiedenen Motorenentwürfen mit ähnlichen Spezifikationen, und so weiter, gesammelt wurden). Der t-Test der gepaarten Zweivariablen ist angemessen, wenn die Vorausberechnungsverteilungen von den exakt gleichen Bevölkerungen stammen (z.B., Daten, die von derselben Kundengruppe, aber bei verschiedenen Gelegenheiten, und so weiter, gesammelt wurden).

Daten extrahieren und Simulationsergebnisse speichern

Man kann die Rohdaten einer Simulation unter Verwendung der *Datenextrahierungsroutine* von Risiko Simulator sehr leicht extrahieren. Man kann sowohl die Hypothesen als auch die Vorausberechnungen extrahieren, aber erst muss eine Simulation ausgeführt werden. Die extrahierten Daten können dann für eine Vielzahl von anderen Analysen verwendet werden.

Prozedur:

- Öffnen Sie oder kreieren Sie ein Modell, definieren Sie Hypothesen und Vorausberechnungen und führen Sie die Simulation aus
- Wählen Sie **Risiko Simulator | Tools | Datenextrahierung**
- Wählen Sie die Hypothesen und/oder Vorausberechnungen, von denen Sie die Daten extrahieren möchten und klicken Sie auf **OK**

Man kann die Daten in verschiedenen Formaten extrahieren:

- Rohdaten in ein neues Arbeitsblatt, wobei man die simulierten Werte (sowohl Hypothesen als auch Vorausberechnungen) speichern oder bei Bedarf weiter analysieren kann
- Flat-Textdatei, wobei man die Daten in andere Datenanalysesoftware exportieren kann
- Risk Simulator Datei, wobei man die Daten (sowohl Hypothesen als auch Vorausberechnungen) zu einem späteren Zeitpunkt abrufen kann, indem man **Risiko Simulator | Tools | Daten öffnen/importieren** auswählt.

Die dritte Option ist die populärste Auswahl. Das heißt, die simulierten Ergebnisse als eine *.risksim Datei speichern, wobei man die Ergebnisse später abrufen kann und eine Simulation nicht jedes Mal ausführen muss. Bild 5.21 zeigt das Dialogfeld für die Extrahierung oder Exportierung und das Speichern der Simulationsergebnisse.

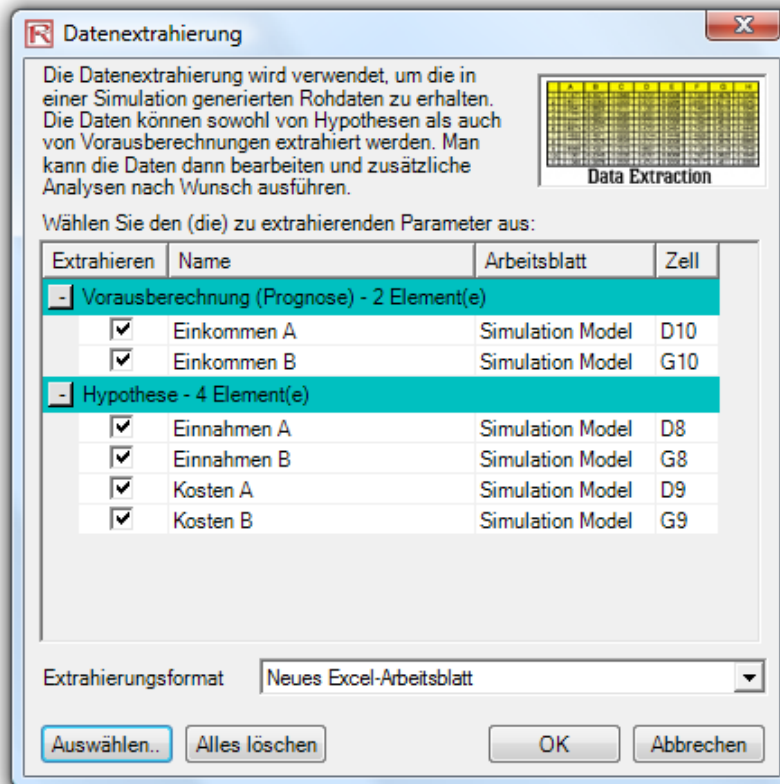


Bild 5.21—Beispiel eines Simulationsberichts

Ein Bericht erstellen

Nach der Ausführung einer Simulation können Sie einen Bericht der Hypothesen, Vorausberechnungen sowie auch der Ergebnisse, die während des Simulationslaufs erhalten wurden, erstellen.

Prozedur:

- Öffnen Sie oder kreieren Sie ein Modell, definieren Sie Hypothesen und Vorausberechnungen und führen Sie die Simulation aus
- Wählen Sie **Risiko Simulator | Bericht kreieren**

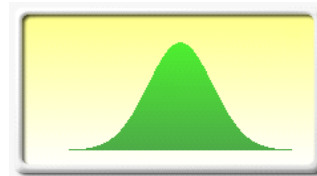
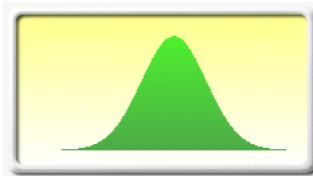
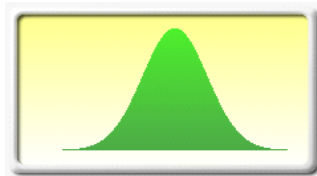
Simulation - Diskontierter Cashflow / ROI Modell

Allgemein

Anzahl der Probeversuche	1000
Simulation bei Fehler an:	Nein
Zufallsausgangswert	123456
Korrelationen aktivieren	Ja

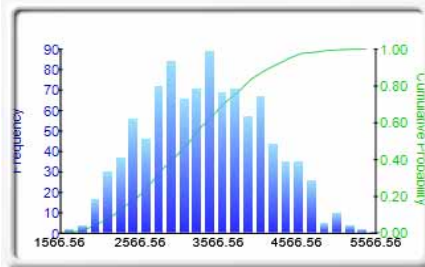
Hypothesen

Name	Product A Avg Price/Unit	Name	Product B Avg Price/Unit	Name	Product C Avg Price/Unit
Aktiviert	Ja	Aktiviert	Ja	Aktiviert	Ja
Zelle	\$C\$12	Zelle	\$C\$13	Zelle	\$C\$14
Dynamische Simulation	Nein	Dynamische Simulation	Nein	Dynamische Simulation	Nein
Bereich		Bereich		Bereich	
Minimum	-Infinity	Minimum	-Infinity	Minimum	-Infinity
Maximum	Infinity	Maximum	Infinity	Maximum	Infinity
Verteilung	Normal	Verteilung	Normal	Verteilung	Normal
Mittelwert	10	Mittelwert	12.25	Mittelwert	15.15
Standardabweichung	1	Standardabweichung	1.225	Standardabweichung	1.515

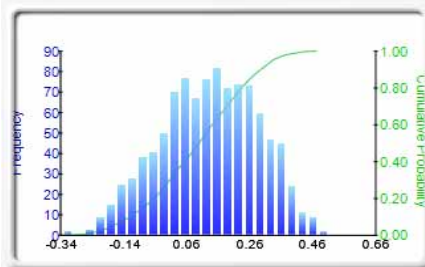


Vorausberechnungen (Prognosen)

Name	Net Present Value	Anzahl der Datenpunkte	1000
Aktiviert	Ja	Mittelwert	3221.5390
Zelle	\$G\$6	Medianwert	3204.1955
Vorausberechnungspräzision		Standardabweichung	744.0429
Präzisionsniveau	---	Varianz	553599.8476
Fehlerniveau	---	Variationskoeffizient	0.2310
		Maximum	5443.1341
		Minimum	1405.0321
		Bereich	4038.1020
		Schiefe	0.1363
		Kurtosis	-0.6015
		25% Perzentil	2663.3412
		75% Perzentil	3764.0708
		Fehlerpräzision bei 95%	0.0143



Name	Intermediate X Variable	Anzahl der Datenpunkte	1000
Aktiviert	Ja	Mittelwert	0.0935
Zelle	\$C\$48	Medianwert	0.0997
Vorausberechnungspräzision		Standardabweichung	0.1549
Präzisionsniveau	---	Varianz	0.0240
Fehlerniveau	---	Variationskoeffizient	1.6570
		Maximum	0.4720
		Minimum	-0.3726
		Bereich	0.8446
		Schiefe	-0.1733
		Kurtosis	-0.5771
		25% Perzentil	-0.0179
		75% Perzentil	0.2111
		Fehlerpräzision bei 95%	0.1027



Korrelationsmatrix

	Variable A	Variable B	Variable C	Variable D	Variable E	Variable F	Variable G	Variable H
Variable A	1.00							
Variable B	0.00	1.00						
Variable C	0.00	0.00	1.00					
Variable D	0.00	0.00	0.00	1.00				
Variable E	0.00	0.00	0.00	0.00	1.00			
Variable F	0.00	0.00	0.00	0.00	0.00	1.00		
Variable G	0.00	0.00	0.00	0.00	0.00	0.00	1.00	
Variable H	0.00	0.00	0.00	0.00	0.00	0.00	0.00	1.00

Bild 5.21—Beispiel eines Simulationsberichts

Regressions- und Vorausberechnungs-Diagnosetool

Dieses fortgeschrittene analytische Tool in Risiko Simulator wird verwendet, um die ökonometrischen Eigenschaften Ihrer Daten zu bestimmen. Die Diagnostik schließt Folgendes ein: Überprüfung der Daten nach Heteroskedastizität, Nichtlinearität, Ausreißer, Spezifikationsfehler, Mikronumerosity, Stationarität und stochastische Eigenschaften, Normalität und Zirkularität der Fehler, und Multikollinearität. Jeder Test wird ausführlicher in dem entsprechenden Bericht beschrieben.

Prozedur:

- Das Beispielsmodell (**Risiko Simulator | Beispiele | Regressionsdiagnostik**) öffnen und zum Arbeitsblatt *Zeitreihendaten* gehen. Wählen Sie die Daten einschließlich der Variablenamen (Zellen **C5:H55**) aus
- Klicken Sie auf **Risiko Simulator | Tools | Diagnosetool**.
- Überprüfen Sie die Daten und wählen Sie die *abhängige Variable Y* aus dem Dropdownmenü aus. Klicken Sie auf **OK**, wenn Sie fertig sind (Bild 5.22).

Multiple Regression Analysis Data Set

Dependent Variable Y	Variable X1	Variable X2	Variable X3	Variable X4	Variable 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	7.2
614	100484	432	29.76	190.8	31.8
385	16728	290	3.294	678.4	340.8
286	14630	346	3.287	239.6	111.9
397	4008	328	0.666	12.938	6.478
764	38927	354	12.938	239.6	111.9
427	22322	266	6.478	111.9	6.478

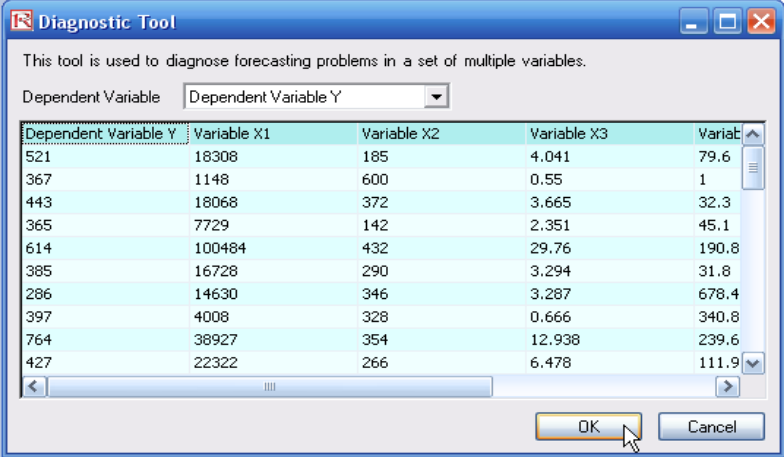


Bild 5.22—Das Datendiagnosetool ausführen

Eine häufige Verletzung bei der Vorausberechnung und der Regressionsanalyse ist die Heteroskedastizität, das heißt, die Varianz der Fehler wächst mit der Zeit (siehe Bild 5.23 für Testergebnisse nach der Verwendung des Diagnosetools). Optisch wachsen oder fächern die vertikalen Datenfluktuationen mit der Zeit auf, und der Bestimmungskoeffizient (R-Quadrat

Koeffizient) fällt typischerweise erheblich, wenn eine Heteroskedastizität vorhanden ist. Wenn die Varianz der abhängigen Variablen nicht konstant ist, dann ist die Varianz des Fehlers nicht konstant. Sofern die Heteroskedastizität der abhängigen Variablen nicht ausgeprägt ist, wird ihre Auswirkung nicht stark sein. Die Schätzungen der kleinsten Quadrate sind trotzdem erwartungstreu und die Schätzungen der Steigung und des Achsabschnitts sind entweder normal verteilt, wenn die Fehler normal verteilt sind, oder zumindest asymptotisch normal verteilt (wenn die Anzahl der Datenpunkte groß wird), wenn die Fehler nicht normal verteilt sind. Die Schätzung der Steigungsvarianz und der Gesamtvarianz werden ungenau sein, aber diese Ungenauigkeit wird wahrscheinlich nicht erheblich sein, wenn die Werte der unabhängigen Variablen symmetrisch um ihren Mittelwert sind.

Wenn die Anzahl der Datenpunkte klein ist (Micronumerosity), könnte es schwer sein, die Hypothesenverletzungen zu erkennen. Bei kleinen Stichproben, sind Hypothesenverletzungen, wie Nichtnormalität oder Heteroskedastizität der Varianzen, schwer aufspürbar, auch wenn sie vorhanden sind. Bei einer kleinen Anzahl von Datenpunkten bietet eine lineare Regression weniger Schutz vor Hypothesenverletzungen. Bei wenigen Datenpunkten könnte es schwer festzustellen sein, wie gut die angepasste Linie den Daten entspricht, oder ob eine nicht lineare Funktion geeigneter wäre. Auch wenn keine der Testhypothesen verletzt wurde, könnte eine lineare Regression auf einer kleinen Anzahl von Datenpunkten beruhend nicht genügend Schärfe besitzen, um einen bedeutenden Unterschied zwischen der Steigung und der Null festzustellen, auch wenn die Steigung ungleich Null ist. Die Schärfe hängt vom Residualfehler, von der beobachteten Variation in der unabhängigen Variablen, vom ausgewählten Alpha-Signifikanzniveau des Tests und von der Anzahl der Datenpunkte ab. Die Schärfe fällt mit dem Steigen der Residualvarianz, fällt mit dem Fallen des Signifikanzniveaus (das heißt, bei einem stringenteren Test), steigt mit dem Ansteigen der Variation in der beobachteten unabhängigen Variablen und steigt mit dem Ansteigen der Datenpunkte.

Die Werte könnten wegen des Vorhandenseins von Ausreißern nicht identisch verteilt sein. Ausreißer sind Abnormwerte in den Daten. Ausreißer können einen starken Einfluss auf die angepasste Steigung und den angepassten Achsabschnitt ausüben, was zu einer schlechten Anpassung des Großteils der Daten führt. Ausreißer neigen dazu, die Schätzung der Residualvarianz zu steigern, was die Chance der Ablehnung der Nullhypothese verringert, sprich, sie erzeugen höhere Vorausberechnungsfehler. Ausreißer können auf Erfassungsfehler zurückzuführen sein, was korrigierbar sein kann, oder sie können auf die Tatsache zurückzuführen sein, dass die Werte der abhängigen Variablen nicht von derselben Bevölkerung entnommen wurden. Scheinausreißer könnten auch auf die Tatsache zurückzuführen sein, dass die Werte der abhängigen Variablen von derselben, aber von einer nicht normalen, Bevölkerung stammen. Allerdings kann ein Punkt ein ungewöhnlicher Wert entweder in einer unabhängigen oder einer abhängigen Variablen sein, ohne unbedingt ein Ausreißer im Streudiagramm sein zu müssen. In einer Regressionsanalyse kann die angepasste Linie höchstanfällig auf Ausreißer sein. In anderen Worten, die Regression der kleinsten Quadrate ist nicht gegen Ausreißer resistent und

demnach ist es die Schätzung der angepassten Steigung auch nicht. Ein von anderen Punkten vertikal entfernter Punkt kann bewirken, dass die angepasste Linie nah an ihm vorbeizieht, anstatt dem allgemeinen Lineartrend der Restdaten zu folgen, insbesondere wenn der Punkt ziemlich weit horizontal von dem Datenzentrum liegt.

Allerdings sollte man bei der Entscheidung, ob die Ausreißer entfernt werden sollen, große Sorgfalt walten lassen. Auch wenn in den meisten Fällen die Regressionsergebnisse besser aussehen, wenn die Ausreißer entfernt werden, muss erst eine *a priori* Rechtfertigung vorhanden sein. Wenn man zum Beispiel eine Regression auf die Wertentwicklung der Aktienrenditen einer gewissen Firma durchführt, sollte man Ausreißer, die von Aktienmarktabschwüngen verursacht wurden, einschließen; dies sind eigentlich keine Ausreißer, da sie Unvermeidbarkeiten des Geschäftszykluses sind. Wenn man auf diese Ausreißer verzichtet und die Regressionsgleichung verwendet, um seinen Pensionsfond auf Basis der Firmenaktien vorauszuberechnen, wird man bestenfalls falsche Ergebnisse bekommen. Nehmen wir im Gegensatz dazu an, dass die Ausreißer von einer einzelnen einmaligen Geschäftslage verursacht wurden (z.B., Fusion und Übernahme) und dass solche Änderungen in der Geschäftsstruktur nicht wieder vorgesehen sind, dann sollte man diese Ausreißer entfernen und die Daten bereinigen, bevor man eine Regressionsanalyse ausführt. Die hier vorgeführte Analyse dient nur zur Identifizierung von Ausreißern. Ob die Ausreißer entfernt werden sollten oder nicht, wird dem Benutzer überlassen.

Gelegentlich ist ein nicht lineares Verhältnis zwischen den abhängigen und unabhängigen Variablen angemessener als ein lineares Verhältnis. In solchen Fällen ist die Ausführung einer linearen Regression nicht optimal. Wenn das lineare Modell nicht die korrekte Form besitzt, dann werden sowohl die Schätzungen der Steigung und des Achsabschnitts als auch die angepassten Werte von der Regression verzerrt sein, und die angepassten Schätzungen der Steigung und des Achsabschnitts werden keine Bedeutung haben. Über einen begrenzten Bereich von unabhängigen oder abhängigen Variablen, können lineare Modelle nicht lineare Modelle gut approximieren (dies ist im Grunde die Basis der linearen Interpolation), aber für eine genaue Vorhersage sollte man ein den Daten angemessenes Modell auswählen. Man sollte zuerst eine nichtlineare Transformation auf die Daten anwenden, bevor man eine Regression durchführt. Eine einfache Methode ist den natürlichen Logarithmus der unabhängigen Variablen zu nehmen (andere Methoden schließen ein, die Quadratwurzel zu ziehen oder die unabhängige Variable in die zweite oder dritte Potenz zu erheben) und eine Regression oder Vorausberechnung unter Verwendung der nichtlinear transformierten Daten durchzuführen.

Ergebnisse der Diagnostik

Variable	Ergebnisse der Diagnostik		Micronumerosity	Ausreißer		Anzahl der potentiellen Ausreißer	Nichtlinearität	
	W-Test p-Wert	Ergebnis des Hypothesentests	Ergebnis der Approximation	Natürliche untere Grenze	Natürliche obere Grenze		Nichtlinear-Tes p-Wert	Ergebnis des Hypothesentests
Y			keine probleme	-7.86	671.70	2		
X1	0.2543	Homoskedastic	keine probleme	-21377.95	64713.03	3	0.2458	linéaire
X2	0.3371	Homoskedastic	keine probleme	77.47	445.93	2	0.0335	non linéaire
X3	0.3649	Homoskedastic	keine probleme	-5.77	15.69	3	0.0305	non linéaire
X4	0.3066	Homoskedastic	keine probleme	-295.96	628.21	4	0.9298	linéaire
X5	0.2495	Homoskedastic	keine probleme	3.35	9.38	3	0.2727	linéaire

Bild 5.23—Ergebnisse von Tests für Ausreißer, Heteroskedastizität, Mikronumerosity und Nichtlinearität

Wenn man Zeitreihendaten vorausberechnet, ist ein typisches Thema, ob die Werte der unabhängigen Variablen wirklich unabhängig oder ob sie doch abhängig von einander sind. Die Werte von abhängigen Variablen, die von einer Zeitreihe gesammelt wurden können autokorreliert sein. Für seriell korrelierte Werte von abhängigen Variablen werden die Schätzungen der Steigung und des Achsabschnitts erwartungstreu sein. Aber die Schätzungen ihrer Vorausberechnungen und Varianzen werden nicht zuverlässig sein und damit wird die Gültigkeit von bestimmten statistischen „Güte-der-Anpassung“ Tests fehlerhaft sein. Zum Beispiel sind Zinssätze, Inflationsraten, Umsätze, Einnahmen und viele andere Zeitreihendaten autokorreliert, wobei der Wert der aktuellen Periode in Zusammenhang mit dem Wert einer früheren Periode steht und so weiter (so ist die Inflationsrate von März mit dem Niveau von Februar verbunden, das seinerseits mit dem Niveau von Januar in Zusammenhang steht und so weiter). Die Nichtbeachtung solcher offensichtlichen Zusammenhänge wird verzerrte und weniger genauere Vorausberechnungen ergeben. In solchen Fällen könnte ein autokorreliertes Regressionsmodell oder ein ARIMA Modell geeigneter sein (**Risiko Simulator I Vorausberechnung I ARIMA**). Zuletzt noch, die Autokorrelationsfunktionen von einer nichtstationären Reihe neigen dazu, langsam zu zerfallen (siehe den Bericht über Nichtstationarität).

Wenn die Autokorrelation $AC(1)$ ungleich Null ist, bedeutet das, dass die Reihe in der 1. Ordnung seriell korreliert ist. Wenn die $AC(k)$ mehr oder weniger geometrisch mit zunehmender Verzögerung wegstirbt, bedeutet das, dass die Reihe einem autoregressiven Prozess der unteren Ordnung folgt. Wenn die $AC(k)$ nach einer kleinen Anzahl von Verzögerungen auf Null fällt, bedeutet das, dass die Reihe einem Prozess mit gleitendem Mittelwert der unteren Ordnung folgt. Die partielle Korrelation $PAC(k)$ misst die Korrelation der Werte, die k Perioden auseinander sind nach Entfernung der Korrelation von den Zwischenverzögerungen. Wenn die Struktur der Autokorrelation von einer Autoregression mit einer Ordnung kleiner als k erfasst werden kann, dann wird die partielle Autokorrelation bei der Verzögerung k nahe Null liegen. Ljung-Box Q-Statistiken und deren p-Werte bei der Verzögerung k weisen die Nullhypothese auf, dass es keine Autokorrelation bis hin zur Ordnung k gibt. Die punktierten Linien in den Diagrammen der Autokorrelationen sind die angenäherten zwei Standardfehlergrenzen. Wenn sich die Autokorrelation innerhalb dieser Grenzen befindet, ist sie nicht signifikant abweichend von Null bei einem 5% Signifikanzniveau.

Die Autokorrelation misst das Verhältnis der Vergangenheit der abhängigen Y-Variablen zu sich selbst. Verteilungsverzögerungen, dagegen, sind Zeitverzögerungsverhältnisse zwischen der abhängigen Y-Variablen und verschiedenen unabhängigen X-Variablen. Zum Beispiel neigen die Bewegung und Richtung von Hypothekensätzen dazu, dem Bundesmittelsatz zu folgen, aber mit einer Zeitverzögerung (typisch sind 1 bis 3 Monate). Zeitverzögerungen folgen gelegentlich Zyklen und Saisonalität (z.B., neigt der Verkauf von Speiseeis dazu, seinen Höhepunkt in den Sommermonaten zu erreichen und stehen deshalb in Zusammenhang mit dem Umsatz des letzten Sommers, 12 Monate in der Vergangenheit). Die Verteilungsverzögerungsanalyse (Bild 5.24) zeigt, wie die abhängige Variable im Zusammenhang mit jeder der unabhängigen Variablen bei verschiedenen Zeitverzögerungen steht, wenn alle Verzögerungen gleichzeitig in Betracht gezogen werden, um festzustellen, welche Zeitverzögerungen statistisch signifikant sind und in Betracht gezogen werden sollten.

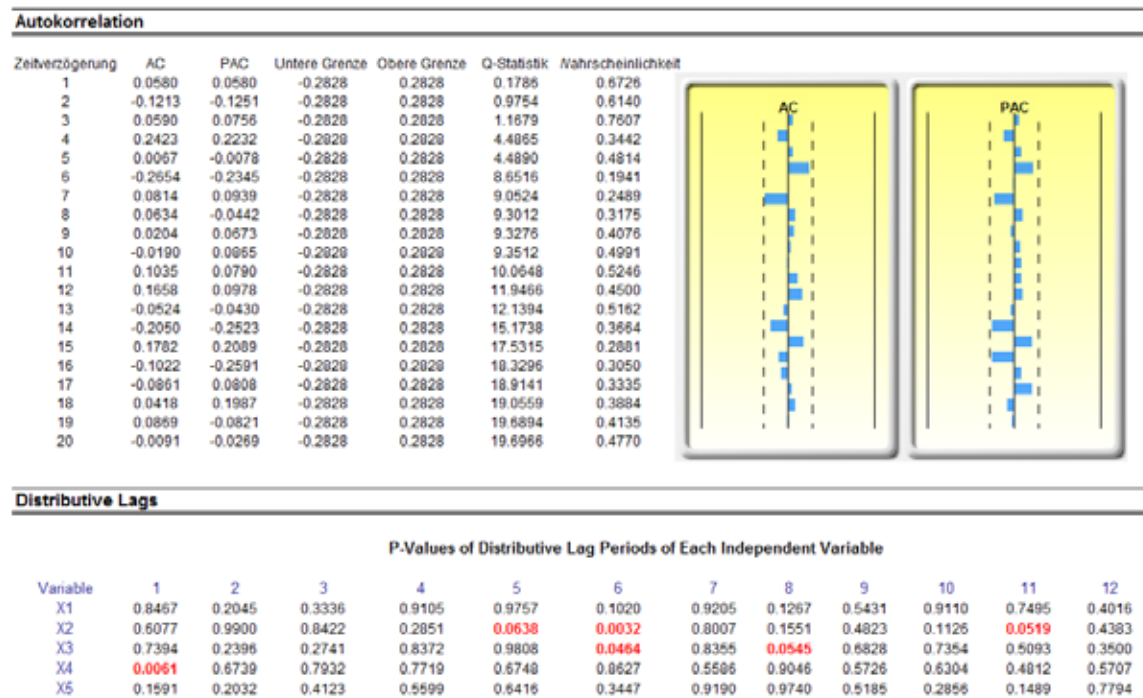


Bild 5.24—Ergebnisse der Autokorrelation und Verteilungsverzögerung

Eine weitere Voraussetzung bei der Ausführung eines Regressionsmodells ist die Annahme der Normalität und Zirkularität der Fehlerbezeichnung. Wenn die Normalitätsannahme verletzt wird oder Ausreißer vorliegen, dann kann die lineare Regression Güte-der-Anpassung Test nicht der schärfste oder informationsreichste der verfügbaren Tests sein, und dies kann den Unterschied zwischen der Feststellung einer linearen Anpassung oder nicht bedeuten. Wenn die Fehler nicht unabhängig und nicht normalverteilt sind, kann es darauf hinweisen, dass die Daten autokorreliert sein können oder unter Nichtlinearitäten oder anderen destruktiver Fehlern leiden. Man kann die Unabhängigkeit der Fehler auch im Heteroskedastizität Test feststellen (Bild 5.25).

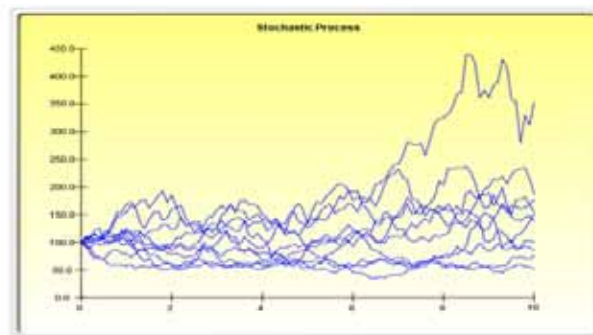
Der durchgeführte Normalitätstest auf Fehler ist ein nicht parametrischer Test, der keine Annahmen macht über die spezifische Form der Bevölkerung aus der die Stichprobe entnommen wurde, was die Analyse von kleineren Datensätzen erlaubt. Dieser Test wertet die Nullhypothese, dass die Stichprobenfehler von einer normalverteilten Bevölkerung stammen, gegen eine Alternativhypothese, dass die Stichprobenfehler von einer nicht normalverteilten Bevölkerung stammen, aus. Wenn die berechnete D-Statistik größer als oder gleich den D-Kritischen Werten bei verschiedenen Signifikanzwerten ist, dann sollte man die Nullhypothese ablehnen und die Alternativhypothese akzeptieren (die Fehler sind nicht normalverteilt). Wenn die D-Statistik andernfalls geringer als der D-Kritische Wert ist, sollte man die Nullhypothese nicht ablehnen (die Fehler sind normalverteilt). Dieser Test stützt sich auf zwei kumulative Frequenzen: Der ersten, die aus dem Stichprobendatensatz hergeleitet ist, und der zweiten aus einer theoretischen Verteilung, die auf dem Mittelwert und der Standardabweichung der Stichprobendaten basiert.

Testergebnis						
		Die Fehler	Relative Frequenz	Beobachtet	Erwartet	B-E
Durchschnitt des Regressionsfehlers	0.00	-219.04	0.02	0.02	0.0612	-0.0412
Standardabweichung der Fehler	141.83	-202.53	0.02	0.04	0.0766	-0.0366
D-Statistik	0.1036	-186.04	0.02	0.06	0.0948	-0.0348
D-Kritischer Wert bei 1%	0.1138	-174.17	0.02	0.08	0.1097	-0.0297
D-Kritischer Wert bei 5%	0.1225	-162.13	0.02	0.10	0.1265	-0.0265
D-Kritischer Wert bei 10%	0.1458	-161.62	0.02	0.12	0.1272	-0.0072
Nullhypothese: Die Fehler sind normalverteilt.		-160.39	0.02	0.14	0.1291	0.0109
		-145.40	0.02	0.16	0.1526	0.0074
		-138.92	0.02	0.18	0.1637	0.0163
		-133.81	0.02	0.20	0.1727	0.0273
		-120.76	0.02	0.22	0.1973	0.0227
		-120.12	0.02	0.24	0.1985	0.0415
		-113.25	0.02	0.26	0.2123	0.0477

Bild 5.25—Test für die Normalität der Fehler

Gelegentlich kann man bestimmte Typen von Zeitreihendaten mit keiner anderen Methode als einem stochastischen Prozess modellieren, weil die unterliegenden Ereignisse von stochastischer Natur sind. Man kann zum Beispiel Aktienpreise, Zinssätze, Erdölpreise und andere Rohstoffpreise nicht angemessen unter Verwendung eines Regressionsmodells modellieren und vorausberechnen, da diese Variablen höchst ungewiss und schwankungsanfällig sind und sie keiner vordefinierten statischen Verhaltensregel folgen. Mit anderen Worten handelt es sich um einen nichtstationären Prozess. Die Stationarität wird hier unter Verwendung des Durchläufe-Tests geprüft, während man einen weiteren visuellen Hinweis in dem Autokorrelationsbericht findet (die Autokorrelationsfunktion neigt dazu, langsam zu zerfallen). Ein stochastischer Prozess ist eine Folge von Ereignissen oder Pfaden, die mittels probabilistischen Gesetzen generiert wurden. Das heißt, Zufallsereignisse können im Laufe der Zeit stattfinden, aber sie werden von spezifischen statistischen und probabilistischen Regeln beherrscht. Die wichtigsten stochastischen Prozesse sind, unter anderem, die Irrfahrt oder Brownsche Bewegung, die Rückkehr zum Mittelwert und die Sprung-Diffusion. Man kann diese Prozesse verwenden, um eine Vielzahl von Variablen, die scheinbar Zufallstrends folgen, aber doch durch probabilistische Gesetze eingeschränkt sind, vorauszuberechnen. Die prozessgenerierende Gleichung ist im Voraus bekannt, aber die tatsächlich generierten Ergebnisse sind unbekannt (Bild 5.26).

Der Prozess mit Irrfahrt Brownsche Bewegung kann verwendet werden, um Aktienpreise, Rohstoffpreise und andere stochastische Zeitreihendaten vorauszuberechnen, gegeben eine Drift oder Wachstumsrate und eine Volatilität um den Driftpfad. Der Prozess mit Rückkehr zum Mittelwert kann verwendet werden, um die Fluktuationen des Irrfahrtprozesses zu reduzieren, indem er dem Pfad erlaubt einen Langzeitwert anzupeilen. Dies macht ihn nützlich, um Zeitreihenvariablen vorauszuberechnen, die eine Langzeitrage haben, wie Zinssätze und Inflationsraten (dies sind Langzeitzielraten der Regulierungsbehörden oder des Markts). Der Prozess mit Sprung-Diffusion ist nützlich, um Zeitreihendaten vorauszuberechnen, wenn die Variable gelegentlich Zufallssprünge aufweisen kann, wie Erdöl- oder Strompreise (diskrete exogene Ereignisschocks können Preise nach oben oder nach unten springen lassen). Letztlich, man kann diese drei stochastischen Prozesse nach Wunsch mischen und anpassen.



Statistische Zusammenfassung

Folgend sind die geschätzten Parameter für einen stochastischen Prozess, gegeben durch die bereitgestellten Daten. Es wird Ihnen überlassen, festzustellen, ob die Wahrscheinlichkeit der Anpassung (ähnlich einer Güte-der-Anpassung Berechnung) reicht, um die Verwendung einer Vorausberechnung mit stochastischem Prozess zu rechtfertigen, und wenn ja, ob man einen Modell mit Irrfahrt, mit Rückkehr zum Mittelwert, mit Sprung-Diffusion oder eine Kombination von diesen verwenden sollte. Bei der Auswahl des richtigen Stochastischem-Prozess-Modells, müssen Sie sich auf frühere Erfahrungen und „a priori“ ökonomische und finanzielle Erwartungen stützen, über was den unterliegende Datensatz am besten repräsentieren kann. Diese Parameter können in eine Vorausberechnung mit stochastischem Prozess eingegeben werden (Risiko Simulator I Vorausberechnung I Stochastische Prozesse).

Periodisch					
Drift rate	-1.48%	Rückkehrsatz	283.89%	Sprungsatz	20.41%
Volatilität	88.84%	Langzeitwert	327.72	Sprunggröße	237.89

Wahrscheinlichkeit der Anpassung eines stochastischen Modells:
Eine hohe Anpassung bedeutet, dass ein stochastisches Modell besser ist als konventionelle Modelle.

Durchläufe	20	Standard Normal	-1.7321
Positive	25	P-Wert (2-Schwanz)	0.0416
Negative	25	P-Value (2-tail)	0.0833
Erwartungsdurchlauf	26		

Ein niedriger p-Wert (unter 0.10, 0.05, 0.01) bedeutet, dass die Sequenz nicht zufällig ist und sie deshalb an Stationaritätsproblemen leidet und ein ARIMA Modell geeigneter sein könnte. In Gegensatz dazu, weisen höhere p-Werte auf eine Zufälligkeit hin und Modelle mit stochastischem Prozess könnten geeignet sein.

Bild 5.26—Stochastischer Prozess - Parameterschätzung

Ein Wort der Warnung an dieser Stelle. Die Kalibrierung der stochastischen Parameter zeigt alle Parameter für alle Prozesse und unterscheidet nicht, welcher Prozess besser und welcher schlechter ist oder welcher Prozess zur Verwendung angemessener ist. Es wird dem Benutzer überlassen, dies zu bestimmen. Zum Beispiel, wenn wir eine Rückkehrrate von 283% sehen, ist aller Wahrscheinlichkeit nach ein Prozess mit Rückkehr zum Mittelwert nicht angemessen. Oder, eine hohe Sprungrate, sagen wir, von 100% bedeutet höchstwahrscheinlich, dass ein Prozess mit Sprung-Diffusion vermutlich nicht angemessen ist, und so weiter. Außerdem kann die Analyse

nicht feststellen, was die Variable ist und was die Datenquelle ist. Zum Beispiel, stammen die Rohdaten von historischer Aktienpreise oder sind sie historische Strompreise oder Inflationsraten oder die molekulare Bewegung von subatomaren Teilchen, und so weiter. Nur der Benutzer würde das wissen und deshalb, unter Verwendung „a priori“ Kenntnis und Theorie, imstande sein, das richtige zu verwendende Verfahren auszuwählen (z.B., Aktienpreise neigen dazu, eine Irrfahrt mit Brownscher Bewegung zu folgen, während Inflationsraten einen Prozess mit Rückkehr zum Mittelwert folgen, aber ein Prozess mit Sprung-Diffusion ist angemessener, wenn Sie den Strompreis vorausberechnen wollen).

Eine Multikollinearität besteht, wenn ein lineares Verhältnis zwischen den unabhängigen Variablen vorhanden ist. Wenn sich das ereignet, kann die Regressionsgleichung gar nicht geschätzt werden. In fast Kollinearitäts-Situationen wird die geschätzte Regressionsgleichung verzerrt sein und ungenaue Ergebnisse liefern. Diese Situation trifft besonders zu, wenn man die Methode der schrittweisen Regression verwendet, wobei die statistisch signifikanten unabhängigen Variablen früher als erwartet aus dem Regressionsmix ausgeworfen werden. Dies ergibt eine Regressionsgleichung, die weder effizient noch genau ist. Ein Schnelltest zur Präsenz einer Multikollinearität in einer mehrfachen Regressionsgleichung ist, dass der R-Quadrat Wert relativ hoch ist, während die t-Statistiken relativ niedrig sind.

Ein weiterer Schnelltest ist es, eine Korrelationsmatrix zwischen den unabhängigen Variablen zu kreieren. Eine hohe Kreuzkorrelation deutet auf das Potential für eine Autokorrelation hin. Eine Faustregel ist, dass eine Korrelation mit einem absoluten Wert höher als 0,75 auf eine schwere Multikollinearität hindeutet. Ein anderer Test für die Multikollinearität ist die Verwendung des Varianzinflationsfaktors (VIF): Man führt eine Regression jeder unabhängigen Variablen zu allen anderen unabhängigen Variablen aus, um den R-Quadrat Wert zu erhalten und somit den VIF zu berechnen. Einen VIF über 2.0 gilt als eine schwere Multikollinearität. Einen VIF über 10.0 deutet auf eine destruktive Multikollinearität hin (Bild 5.27).

Korrelationsmatrix					
KORRELATION	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	
				1.000	

Varianzinflationsfaktor					
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	
				N/A	

Bild 5.27—Multikollinearitätsfehler

Die Korrelationsmatrix listet die Produkt-Moment-Korrelationen von Pearson (gewöhnlich Pearsons R genannt) zwischen Variablenpaaren auf. Der Korrelationskoeffizient bewegt sich zwischen -1.0 und + 1.0 inklusiv. Die Zeichen zeigen die Richtung der Assoziation zwischen den Variablen an, während der Koeffizient die Größe oder Stärke der Assoziation anzeigt. Das Pearsons R misst nur ein lineares Verhältnis und ist weniger effektiv bei der Messung von nichtlinearen Verhältnissen.

Um zu testen, ob die Korrelationen signifikant sind, wird ein Zweiseitiger Hypothesentest durchgeführt und die resultierenden p-Werte sind oben aufgelistet. P-Werte kleiner als 0.10, 0.05 und 0.01 sind in Blau hervorgehoben, um eine statistische Signifikanz anzuzeigen. In anderen Worten, einen p-Wert für ein Korrelationspaar der kleiner als ein gegebener Signifikanzwert ist, ist statistisch signifikant abweichend von Null, was darauf hindeutet, dass es ein signifikantes lineares Verhältnis zwischen den beiden Variablen gibt.

Der Koeffizient (R) der Produkt-Moment-Korrelation von Pearson zwischen zwei Variablen (x und y) steht in Zusammenhang zum Kovarianzmaß (cov), wobei $R_{x,y} = \frac{COV_{x,y}}{s_x s_y}$. Der Vorteil des

Teilens der Kovarianz durch das Produkt der Standardabweichung (s) der zwei Variablen ist, dass der resultierende Korrelationskoeffizient zwischen -1.0 und +1.0 inklusiv begrenzt ist. Dies macht die Korrelation zu einem guten relativen Maß, um Vergleiche zwischen verschiedenen Variablen durchzuführen (insbesondere bei unterschiedlichen Maßeinheiten und Größen). Die nicht-parametrische, rangbasierte Korrelation von Spearman wird ebenso unten einbezogen. Das Spearman R steht in Zusammenhang mit dem Pearson R insofern, als die Daten erst per Rang geordnet und dann korreliert werden. Die Rangkorrelationen liefern eine bessere Schätzung der Verhältnisse zwischen zwei Variablen, wenn eine oder beide nicht linear sind.

Man muss betonen, dass eine signifikante Korrelation keine Verursachung impliziert. Assoziationen zwischen Variablen implizieren keineswegs, dass die Änderung einer Variablen eine Änderung in einer anderen Variablen verursacht. Wenn zwei Variable sich unabhängig von einander aber in einem verwandten Pfad bewegen, könnten sie korreliert sein, aber ihr Verhältnis könnte dennoch ein Scheinverhältnis sein (z.B., eine Korrelation zwischen Sonnenflecken und dem Aktienmarkt kann stark sein, aber man kann erraten, dass es keine Kausalität gibt und dass dies ein reines Scheinverhältnis ist).

Statistische Analyse Tool

Ein weiteres sehr leistungsstarkes Tool in Risiko Simulator ist das Statistische Analyse Tool, welches die statistischen Eigenschaften der Daten feststellt. Die ausgeführte Diagnostik schließt das Prüfen der Daten für verschiedene statistische Eigenschaften ein, von elementaren

deskriptiven Statistiken bis hin zum Testen für und Kalibrierung der stochastischen Eigenschaften der Daten.

Prozedur:

- Öffnen Sie das Beispielsmodell (**Risiko Simulator | Beispiele | Statistische Analyse**), gehen Sie zum Arbeitsblatt *Daten* und wählen Sie die Daten einschließlich der Variablenamen aus (Zellen **C5:E55**).
- Klicken Sie auf **Risiko Simulator | Tools | Statistische Analyse** (Bild 5.28).
- Prüfen Sie den *Datentyp*, ob die ausgewählten Daten von einer einzelnen Variable oder von mehrfachen in Reihen geordneten Variablen stammen. In unserem Beispiel, nehmen wir an, dass die ausgewählten Datenbereiche von mehrfachen Variablen stammen. Wenn fertig, klicken Sie auf **OK**
- Wählen Sie die statistischen Tests, die Sie ausführen möchten. Die Empfehlung (und die Standardeinstellung) ist, alle Tests auszuwählen. Wenn fertig, klicken Sie auf **OK** (Bild 5.29).

Nehmen Sie sich Zeit bei der Examinierung der generierten Berichte, um eine bessere Kenntnis der ausgeführten statistischen Tests zu bekommen (Beispielsberichte werden in Bilder 5.30-5.33 angezeigt).

Datensatz

Variable X1	Variable X2	Variable X3
521	18308	185
367	1148	600
443	18068	372
365	7729	142
614	100484	432
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	

Statistische Analyse

Dieses Tool wird verwendet, um die statistischen Verhältnisse in einem Satz von Rohdaten zu beschreiben und zu finden.

Ausgewählte Daten

Variable X1	Variable X2	Variable X3
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

☐ Die Daten stammen von einer einzelnen Variablen.
 ☒ Die Daten umfassen mehrfache Variablen in Spalten.

OK Abbrechen

Bild 5.28—Das Statistische Analyse Tool ausführen

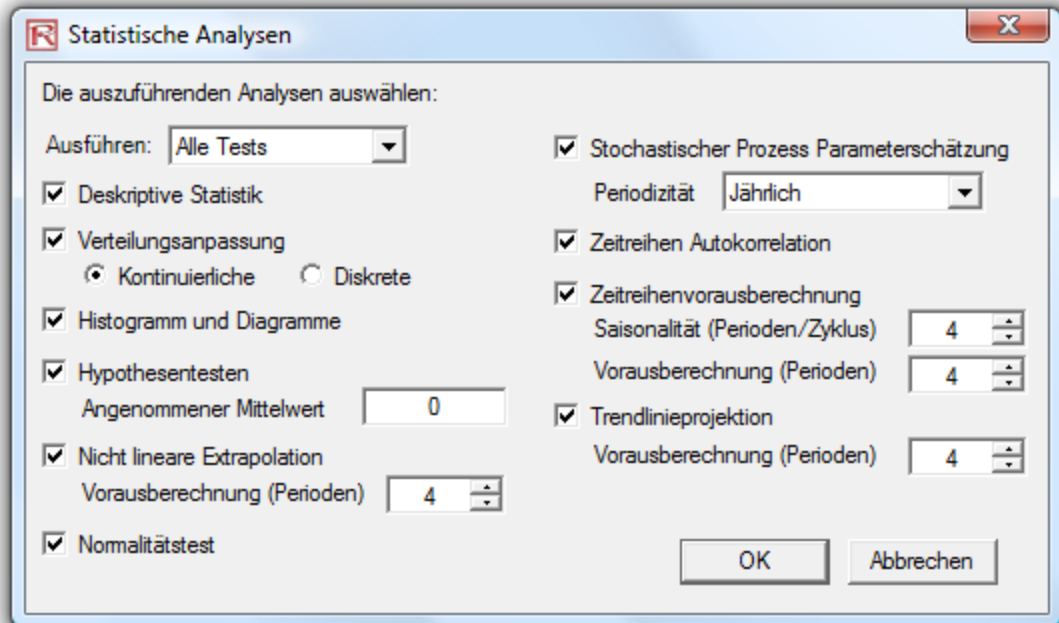


Bild 5.29—Statistische Tests

Deskriptive Statistik

Analyse von Statistiken

Beinah alle Verteilungen können innerhalb von 4 Momenten beschrieben werden (einige Verteilungen erfordern einen Moment, während andere zwei Momente erfordern und so weiter). Die deskriptiven Statistiken erfassen quantitativ diese Momente. Der erste Moment beschreibt die Position einer Verteilung (sprich, Mittelwert, Medianwert und Modalwert) und wird als der Erwartungswert, die Erwartungsrenditen oder der Durchschnittswert der Ereignisse gedeutet.

Der arithmetische Mittelwert berechnet den Durchschnitt aller Ereignisse, indem alle Datenpunkte addiert und durch die Anzahl der Punkte geteilt werden. Der geometrische Mittelwert wird berechnet, indem man die Potenzwurzel der Produkte aller Datenpunkte zieht, die alle positiv sein müssen. Der geometrische Mittelwert ist genauer für Prozentsätze und Sätze, die erheblich fluktuieren. Sie können den geometrischen Mittelwert zum Beispiel verwenden, um die Durchschnittswachstumsrate zu berechnen, gegeben den Zinssatz mit variablen Sätzen. Der getrimmte Mittelwert berechnet den arithmetischen Durchschnitt des Datensatzes, nachdem die extremen Ausreißer getrimmt wurden. Da Durchschnittswerte bei Anwesenheit von Ausreißern zu einer signifikanten Verzerrung neigen, reduziert der getrimmte Mittelwert solche Verzerrungen in asymmetrischen Verteilungen.

Der Standardfehler des Mittelwerts berechnet den Fehler um den Stichprobenmittelwert. Je größer die Stichprobengröße, umso kleiner ist der Fehler, so dass für eine unendlich große Stichprobe sich der Fehler der Null nähert. Dies zeigt an, dass der Bevölkerungsparameter geschätzt wurde. Auf Grund der Stichprobenentnahmefehler, wird der 95% Konfidenzintervall für den Mittelwert bereitgestellt. Basierend auf einer Analyse der Stichprobendatenpunkte, sollte der Bevölkerungsmittelwert zwischen die unteren und oberen Intervalle für den Mittelwert fallen.

Der Medianwert ist der Datenpunkt bei dem 50% aller Datenpunkte über diesem Wert und 50% unter diesen Wert liegen. Unter diesen drei ersten Statistikmomenten ist der Medianwert am wenigsten anfällig auf Ausreißer. Eine symmetrische Verteilung besitzt einen Medianwert, der dem arithmetischen Mittelwert gleicht. Eine asymmetrische Verteilung existiert, wenn der Medianwert weit entfernt vom Mittelwert liegt. Der Modalwert misst den am häufigsten auftretenden Datenpunkt.

Das Minimum ist der kleinste Wert im Datensatz, während das Maximum der größte Wert ist. Die Schwankungsbreite ist die Abweichung zwischen den Maximal- und Minimalwerten.

Der zweite Moment misst die Streuung oder Breite einer Verteilung und wird oft unter Verwendung von Maßen wie Standardabweichungen, Varianzen, Quartile und Interquartilsabständen beschrieben. Die Standardabweichung zeigt die Durchschnittsabweichung aller Datenpunkte von ihrem Mittelwert. Dies ist eine häufige Maßeinheit, da sie mit dem Risiko assoziiert ist (höhere Standardabweichungen bedeuten eine breitere Verteilung, ein höheres Risiko oder eine breitere Dispersion der Datenpunkte um den Mittelwert) und ihre Einheiten identisch mit dem Originaldatensatz sind. Die Standardabweichung der Stichprobe unterscheidet sich von der Standardabweichung der Bevölkerung darin, dass die erstere eine Freiheitsgradkorrektur verwendet um kleine Stichprobengrößen z.B. Es werden auch untere und obere Konfidenzintervalle für die Standardabweichung bereitgestellt und die wahre Standardabweichung der Bevölkerung fällt innerhalb dieses Intervalls. Wenn Ihr Datensatz alle Elemente der Bevölkerung abdeckt, verwenden Sie stattdessen die Standardabweichung der Bevölkerung. Die zwei Varianzmaße sind lediglich die Quadratwerte der Standardabweichungen.

Der Variabilitätskoeffizient ist die Standardabweichung der Stichprobe geteilt durch den Stichprobenmittelwert und beweist einen einheitsfreien Dispersionsmaß, der über verschiedene Verteilungen verglichen werden kann (man kann jetzt Wertverteilungen, die in Millionen von Dollar angegeben sind, mit denen, die in Billionen von Dollar angegeben sind, vergleichen, oder Meter und Kilogramme, usw.). Das 1. Quartil misst das 25. Perzentil der Datenpunkte, wenn sie von ihrem kleinsten zum höchsten Wert angeordnet sind. Das 3. Quartil ist der Wert des 75. Perzentil Datenpunkt. Gelegentlich werden Quartile als die oberen und unteren Schwankungsbreiten einer Verteilung verwendet, da dies den Datensatz stützt, um Ausreißer zu ignorieren. Der Interquartilsabstand ist die Differenz zwischen dem 3. und dem 1. Quartil. Er wird oft verwendet, um die Breite der Spannenmitte einer Verteilung zu messen.

Die Schiefe ist der 3. Moment in einer Verteilung. Die Schiefe bezeichnet den Asymmetriegrad einer Verteilung um ihren Mittelwert. Eine positive Schiefe deutet auf eine Verteilung mit einem asymmetrischen Schwanz hin, der sich mehr in Richtung positiver Werte streckt. Eine negative Schiefe deutet auf eine Verteilung mit einem asymmetrischen Schwanz hin, der sich mehr in Richtung negativer Werte streckt.

Die Wölbung bezeichnet die relative Spitze oder Flachheit einer Verteilung im Vergleich zu einer Normalverteilung. Dies ist der 4. Moment in einer Verteilung. Ein positiver Wölbungswert deutet auf eine relativ spitze Verteilung hin. Eine negative Wölbung deutet auf eine relativ flache Verteilung hin. Die hier gemessene Wölbung Kurtosis wurde auf Null zentriert (bestimmte andere Wölbungsmaßeinheiten sind um 3,0 zentriert). Obwohl beide gleich valide sind, ist die Interpretierung bei einer Zentrierung auf Null einfacher. Eine hohe positive Wölbung deutet auf eine Verteilung, die um ihrer Spannenmitte gespitzt ist und auf leptokurtische oder dicke Schwänze hin. Dies deutet auf eine höhere Wahrscheinlichkeit von extremen Ereignissen hin (z.B. katastrophale Ereignisse, Terroranschläge, Börsenstürze) als bei einer Normalverteilung vorhergesagt werden.

Zusammenfassung der Statistiken

Statistiken	Variable X1		
Beobachtungen	50.0000	Standardabweichung (Stichprobe)	172.9140
Arithmetischer Mittelwert	331.9200	Standardabweichung (Bevölkerung)	171.1761
Geometrischer Mittelwert	281.3247	Unteres Konfidenzintervall für die Standardabweichung	148.6090
Getrimmter Mittelwert	325.1739	Oberes Konfidenzintervall für die Standardabweichung	207.7947
Standardfehler des arithmetischen Mittelwerts	24.4537	Varianz (Stichprobe)	29899.2588
Unteres Konfidenzintervall für den Mittelwert	283.0125	Varianz (Bevölkerung)	29301.2736
Oberes Konfidenzintervall für den Mittelwert	380.8275	Variabilitätskoeffizient	0.5210
Medianwert	307.0000	1. Quartil (Q1)	188.0000
Minimum	47.0000	3. Quartil (Q3)	435.0000
Maximum	764.0000	Interquartilsabstand	247.0000
Spannenbreite	717.0000	Schiefe	0.4838
		Wölbung	-0.0952

Bild 5.30—Statistische Analyse Tool - Beispiel eines Berichts

Hypothesentest (t-Test auf den Bevölkerungsmittelwert von einer Variablen)

Statistische Zusammenfassung

Statistiken vom Datensatz:

Beobachtungen 50
Mittelwert der Stichprobe 331.92
Mittelwert der Standardabweichung 172.91

Vom Benutzer bereitgestellten Statistiken:

Angenommener Mittelwert 0.00

Berechnete Statistiken:

t-Statistik 13.5734
P-Wert (Rechtsschwanz) 0.0000
P- Wert (Linksschwanz) 1.0000
P- Wert (Zweischwanz) 0.0000

Nullhypothese (Ho): μ = angenommener Mittelwert
Alternativhypothese (Ha): μ = angenommener Mittelwert
Bemerkung "<>" bezeichnet "größer als" für Rechtsschwanz, "kleiner als" für Linksschwanz oder "nicht gleich wie" für Zweischwanz

Zusammenfassung des Hypothesentests

Der Ein-Variable t-Test ist geeignet, wenn die Standardabweichung der Bevölkerung unbekannt ist, aber die Stichprobenentnahme der Verteilung als ungefähr normal angenommen wird (der t-Test wird verwendet, wenn die Stichprobengröße kleiner als 30 ist, aber er ist auch geeignet und liefert sogar konservativere Ergebnisse bei größeren Datensätzen). Dieser t-Test kann bei drei Typen von Hypothesentests angewendet werden: einem Zweischwanz Test, einem Rechtsschwanz Test und einem Linksschwanz Test. Alle drei Tests und ihre jeweilige Ergebnisse werden für Sie zum Nachschlagen

Zweischwanz Hypothesentest

Eine Zweischwanzhypothese testet die Nullhypothese Ho, so dass der Bevölkerungsmittelwert statistisch identisch mit dem angenommenen Mittelwert ist. Die Alternativhypothese ist, dass der reelle Bevölkerungsmittelwert statistisch abweichend vom angenommenen Mittelwert ist, wenn unter Verwendung des Stichprobendatensatzes getestet wird. Bei der Verwendung des t-Tests, wenn der berechnete p-Wert kleiner als eine angegebene Signifikanzmenge (typisch 0.10, 0.05 oder 0.01) ist, bedeutete es, dass der Bevölkerungsmittelwert statistisch signifikant abweichend von dem angenommenen Mittelwert bei den 10%, 5% und 1% Signifikanzwerten (oder bei 90%, 95% und 99% statistischen Konfidenz) ist. Wenn in Gegensatz dazu der p- Wert größer als 0.10, 0.05 oder 0.01 ist, ist der Bevölkerungsmittelwert statistisch identisch mit dem angenommenen Mittelwert und alle Abweichungen sind auf reinen Zufall zurückzuführen.

Rechtsschwanz Hypothesentest

Eine Rechtsschwanzhypothese testet die Nullhypothese Ho, so dass der Bevölkerungsmittelwert statistisch kleiner als oder gleich dem angenommenen Mittelwert ist. Die Alternativhypothese ist, dass der reelle Bevölkerungsmittelwert statistisch größer als der angenommene Mittelwert ist, wenn unter Verwendung des Stichprobendatensatzes getestet wird. Bei der Verwendung des t-Tests, wenn der berechnete p-Wert kleiner als eine angegebene Signifikanzmenge (typisch 0.10, 0.05 oder 0.01) ist, bedeutete es, dass der Bevölkerungsmittelwert statistisch signifikant größer als der angenommene Mittelwert bei den 10%, 5% und 1% Signifikanzwerten (oder bei 90%, 95% und 99% statistischen Konfidenz) ist. Wenn in Gegensatz dazu der p- Wert größer als 0.10, 0.05 oder 0.01 ist, ist der Bevölkerungsmittelwert statistisch ähnlich dem oder kleiner als der angenommene Mittelwert.

Linksschwanz Hypothesentest

Eine Linksschwanzhypothese testet die Nullhypothese Ho, so dass der Bevölkerungsmittelwert größer als oder gleich dem angenommenen Mittelwert ist. Die Alternativhypothese ist, dass der reelle Bevölkerungsmittelwert statistisch kleiner als der angenommene Mittelwert ist, wenn unter Verwendung des Stichprobendatensatzes getestet wird. Bei der Verwendung des t-Tests, wenn der berechnete p-Wert kleiner als eine angegebene Signifikanzmenge (typisch 0.10, 0.05 oder 0.01) ist, bedeutete es, dass der Bevölkerungsmittelwert statistisch signifikant kleiner als der angenommene Mittelwert bei den 10%, 5% und 1% Signifikanzwerten (oder bei 90%, 95% und 99% statistischen Konfidenz) ist. Wenn in Gegensatz dazu der p- Wert größer als 0.10, 0.05 oder 0.01 ist, ist der Bevölkerungsmittelwert statistisch ähnlich dem oder größer als der angenommene Mittelwert und alle Abweichungen sind auf reinen Zufall zurückzuführen.

Weil der t-Test konservativer ist und nicht eine bekannte Standardabweichung der Bevölkerung erfordert, wie beim Z-Test, verwenden wir nur diesen t-Test.

Bild 5.31—Statistische Analyse Tool - Beispiel eines Berichts (Hypothesentest von einer Variablen)

Normalitätstest

Der Normalitätstest ist eine Form eines nicht-parametrischen Tests, der keine Annahmen macht über die spezifische Form der Bevölkerung aus der die Stichprobe entnommen wurde, was die Analyse von kleineren Datensätzen erlaubt. Dieser Test wertet die Nullhypothese, dass die Stichprobenfehler von einer normalverteilten Bevölkerung entstammen, gegen eine Alternativhypothese, dass die Stichprobenfehler von einer nicht normalverteilten Bevölkerung entstammen, aus. Wenn der berechnete p-Wert kleiner als oder gleich dem Alphasignifikanzwert ist, dann sollte man die Nullhypothese ablehnen und die Alternativhypothese akzeptieren. Wenn der p-Wert andernfalls größer als der Alphasignifikanzwert ist, sollte man die Nullhypothese nicht ablehnen. Dieser Test stützt sich auf zwei kumulative Frequenzen: Der ersten, die aus dem Stichprobendatensatz hergeleitet ist, und der zweiten aus einer theoretischen Verteilung, die auf dem Mittelwert und der Standardabweichung der Stichprobendaten basiert. Eine Alternative zu diesem Test ist der Chi-Quadrat Normalitätstest. Der Chi-Quadrat-Test erfordert mehr Datenpunkte zur Ausführung im Vergleich mit dem hier verwendeten Normalitätstest.

Testergebnis

		Daten	Relative Frequenz	Beobachtet	Erwartet	B-E
Datendurchschnitt	331.92					
Standardabweichung	172.91	47.00	0.02	0.02	0.0497	-0.0297
D-Statistik	0.0859	68.00	0.02	0.04	0.0635	-0.0235
D-Kritischer bei 1%	0.1150	87.00	0.02	0.06	0.0783	-0.0183
D-Kritischer bei 5%	0.1237	96.00	0.02	0.08	0.0862	-0.0062
D-Kritischer bei 10%	0.1473	102.00	0.02	0.10	0.0918	0.0082
Nullhypothese: Die Fehler sind normalverteilt.		108.00	0.02	0.12	0.0977	0.0223
		114.00	0.02	0.14	0.1038	0.0362
		127.00	0.02	0.16	0.1180	0.0420
		153.00	0.02	0.18	0.1504	0.0296
		177.00	0.02	0.20	0.1851	0.0149
		186.00	0.02	0.22	0.1994	0.0206
		188.00	0.02	0.24	0.2026	0.0374
		198.00	0.02	0.26	0.2193	0.0407
		222.00	0.02	0.28	0.2625	0.0175

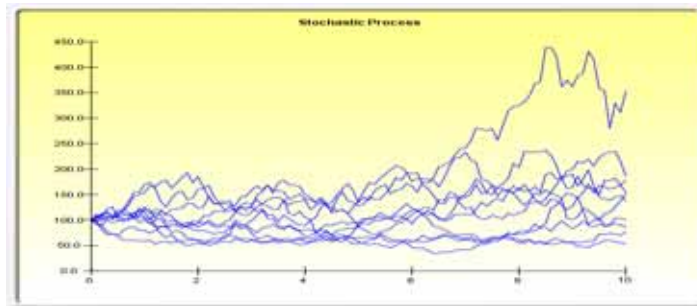
Bild 5.32—Statistische Analyse Tool - Beispiel eines Berichts (Normalitätstest)

Stochastischer Prozess - Parameterschätzungen

Statistische Zusammenfassung

Ein stochastischer Prozess ist eine Folge von Ereignissen oder Pfaden, die mittels probabilistischer Gesetzen generiert wurden. Das heißt, Zufallsereignisse können im Laufe der Zeit stattfinden, aber sie werden von spezifischen statistischen und probabilistischen Regeln beherrscht. Die wichtigsten stochastischen Prozesse sind, unter anderem, die Irrfahrt oder Brownsche Bewegung, die Rückkehr zum Mittelwert und die Sprung-Diffusion. Man kann diese Prozesse verwenden, um eine Vielzahl von Variablen, die scheinbar Zufallstrends folgen aber doch durch probabilistische Gesetze eingeschränkt sind, vorzuberechnen. Die prozessgenerierende Gleichung ist im Voraus bekannt, aber die tatsächlich generierten Ergebnisse sind unbekannt.

Der Prozess mit Irrfahrt Brownsche Bewegung kann verwendet werden, um Aktienpreise, Warenpreise und andere stochastische Zeitreihendaten vorzuberechnen, gegeben eine Drift oder Wachstumsrate und eine Volatilität um den Driftpfad. Der Prozess mit Rückkehr zum Mittelwert kann verwendet werden, um die Fluktuationen des Irrfahrtprozesses zu reduzieren, indem er dem Pfad erlaubt einen Langzeitwert anzupfeilen. Dies macht ihn nützlich, um Zeitreihenvariablen vorzuberechnen, die eine Langzeitrate haben, wie Zinssätze und Inflationsraten (dies sind Langzeitzielen der Regulierungsbehörden oder des Markts). Der Prozess mit Sprung-Diffusion ist nützlich, um Zeitreihendaten vorzuberechnen, wenn die Variable gelegentlich Zufallssprünge aufweisen kann, wie Erdöl- oder Strompreise (diskrete exogene Ereignisschocks können Preise nach oben oder nach unten springen lassen). Letztlich, man kann diese drei stochastischen Prozesse nach Wunsch mischen und anpassen.



Statistische Zusammenfassung

Folgend sind die geschätzten Parameter für einen stochastischen Prozess, gegeben durch die bereitgestellten Daten. Es wird Ihnen überlassen, festzustellen, ob die Wahrscheinlichkeit der Anpassung (ähnlich einer Güte-der-Anpassung Berechnung) reicht, um die Verwendung einer Vorusberechnung mit stochastischem Prozess zu rechtfertigen, und wenn ja, ob man einen Modell mit Irrfahrt, mit Rückkehr zum Mittelwert, mit Sprung-Diffusion oder eine Kombination von diesen verwenden sollte. Bei der Auswahl des richtigen Stochastischen-Prozess-Modells, müssen Sie sich auf frühere Erfahrungen und „a priori“ ökonomische und finanzielle Erwartungen stützen, über was den unterliegende Datensatz am besten repräsentieren kann. Diese Parameter können in eine Vorusberechnung mit stochastischem Prozess eingeben werden (**Risiko Simulator I Vorusberechnung | Stochastische Prozesse**).

(Annualisiert)

Driftrate*	-1.48%	Rückkehrsatz**	283.89%	Sprungsatz**	20.41%
Volatilität*	88.84%	Langzeitwert**	327.72	Sprunggröße**	237.89

Wahrscheinlichkeit der Anpassung eines stochastischen: 46.48%

*Die Werte sind annualisiert

** Die Werte sind periodisch

Bild 5.33—Statistische Analyse Tool - Beispiel eines Berichts (Stochastische Parameterschätzung)

Verteilungsanalyse Tool

Dies ist ein Statistische Wahrscheinlichkeit Tool in Risiko Simulator, das ziemlich nützlich in verschiedenen Bereichen ist. Man kann es bei der Berechnung der Wahrscheinlichkeitsdichtefunktion (PDF) verwenden, welche auch die Wahrscheinlichkeitsfunktion (PMF) genannt wird (wir werden diese Begriffe abwechselnd verwenden), um diskrete Verteilungen zu berechnen, wobei, gegeben eine bestimmte Verteilung und ihre Parameter, wir die Auftretenswahrscheinlichkeit, gegeben einen Ausgang x , feststellen können. Außerdem kann man die kumulative Verteilungsfunktion (CDF) berechnen, welche die der PDF Werte bis zu diesem x Wert ist. Zum Schluss, die inverse kumulative Verteilungsfunktion (ICDF) wird verwendet, um den x Wert, gegeben die kumulative Auftretenswahrscheinlichkeit, zu berechnen.

Dieses Tool ist unter **Risiko Simulator | Tools | Verteilungsanalyse** aufrufbar. Als Beispiel, Bild 5.34 zeigt die Berechnung einer Binomialverteilung (das heißt, eine Verteilung mit zwei Ausgängen, wie zum Beispiel einen Münzwurf, wobei der Ausgang entweder Kopf oder Zahl ist, mit irgendeiner vorgeschriebenen Wahrscheinlichkeit von Köpfen oder Zahlen). Angenommen wir werfen die Münze zwei Male und stellen den Ausgang Kopf als einen Erfolg ein, verwenden wir die Binomialverteilung mit Probeversuchen = 2 (zwei Münzwurfe) und Wahrscheinlichkeit = 0.50 (die Erfolgswahrscheinlichkeit Köpfe zu bekommen). Wenn wir die Wahrscheinlichkeitsdichtefunktion (PDF) auswählen und den Wertebereich x auf 0 bis 2 mit einer Schrittgröße 1 einstellen (dies bedeutet, dass wir die Werte 0, 1, 2 für x erhalten wollen), werden die resultierenden Wahrscheinlichkeiten sowie auch die theoretischen vier Momente der Verteilung in der Tabelle und in einem graphischen Format zur Verfügung gestellt. Da die Ausgänge des Münzwurfs Kopf-Kopf, Zahl-Zahl, Kopf-Zahl und Zahl-Kopf sind, liegt die Wahrscheinlichkeit überhaupt keine Köpfe zu bekommen bei 25%, einen Kopf bei 50% und zwei Köpfe bei 25%.

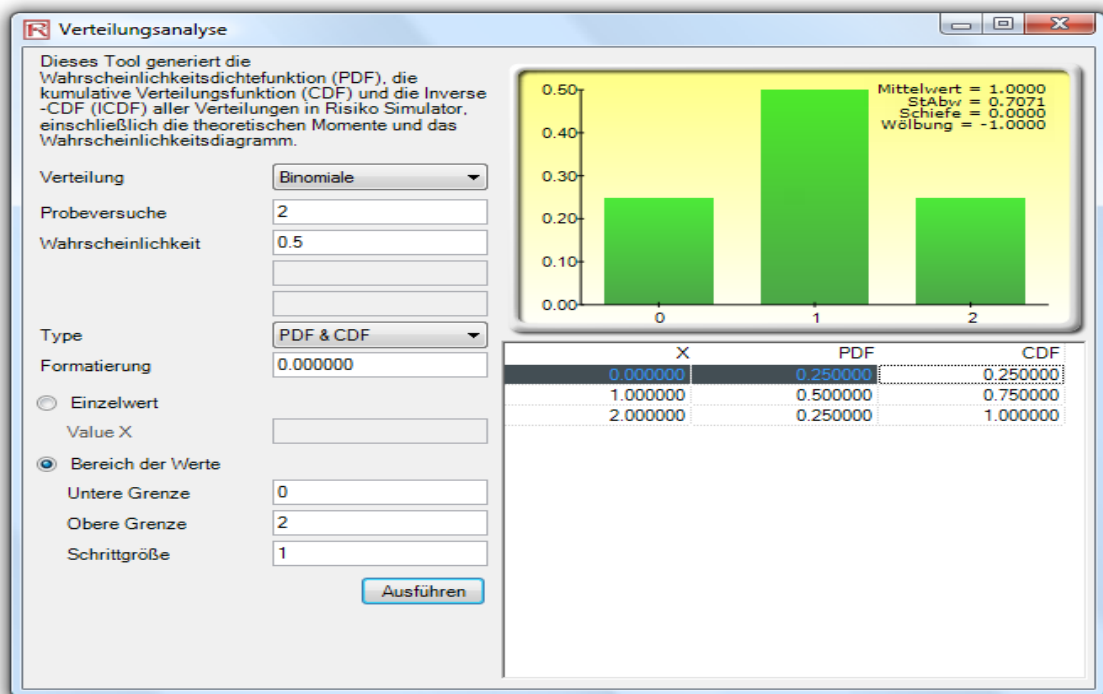


Bild 5.34—Verteilungsanalyse Tool (Binomialverteilung mit 2 Probeversuchen)

Wir können so auch die genauen Wahrscheinlichkeiten von, sagen wir, 20 Münzwürfen bekommen (siehe Bild 5.35). Die Ergebnisse werden in tabellarischen und graphischen Formaten angezeigt.

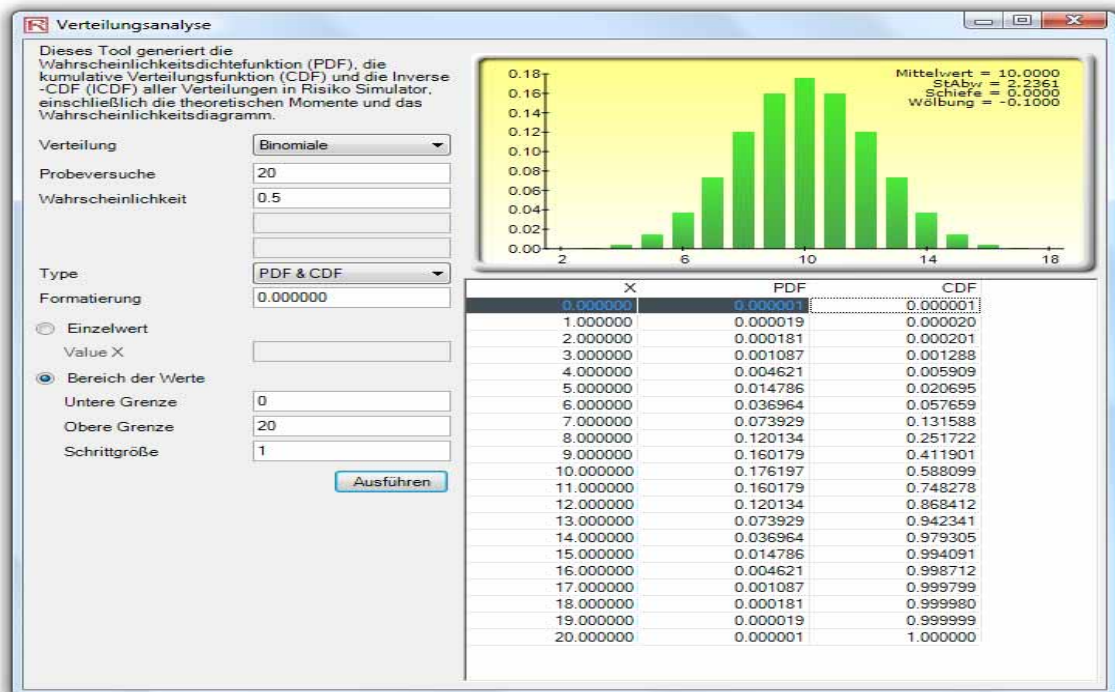


Bild 5.35—Verteilungsanalyse Tool (Binomialverteilung mit 20 Probeversuchen)

Bild 5.36 zeigt die gleiche Binomialverteilung, aber jetzt wird die kumulative Verteilungsfunktion (CDF) berechnet. Die CDF ist lediglich die Summe der Werte der Wahrscheinlichkeitsdichtefunktion (PDF) bis zum Punkt x . Zum Beispiel, im Bild 5.35 sehen wir: Die Wahrscheinlichkeiten von 0, 1 und 2 liegen bei 0.000001, 0.000019 und 0.000181, dessen Summe 0.000201 ist, was der Wert der CDF bei $x = 2$ im Bild 5.36 ist. Während die PDF die Wahrscheinlichkeiten exakt 2 Köpfe zu bekommen berechnet, berechnet die CDF die Wahrscheinlichkeit, dass man nicht mehr als oder bis zu zwei Köpfen bekommen wird (oder Wahrscheinlichkeiten von 0, 1 und 2 Köpfen). Das Komplement (das heißt, $1 - 0.00021$ erhält 0.999799 oder 99.9799%) ist die Wahrscheinlichkeit, dass man mindestens 3 oder mehr Köpfe bekommt.

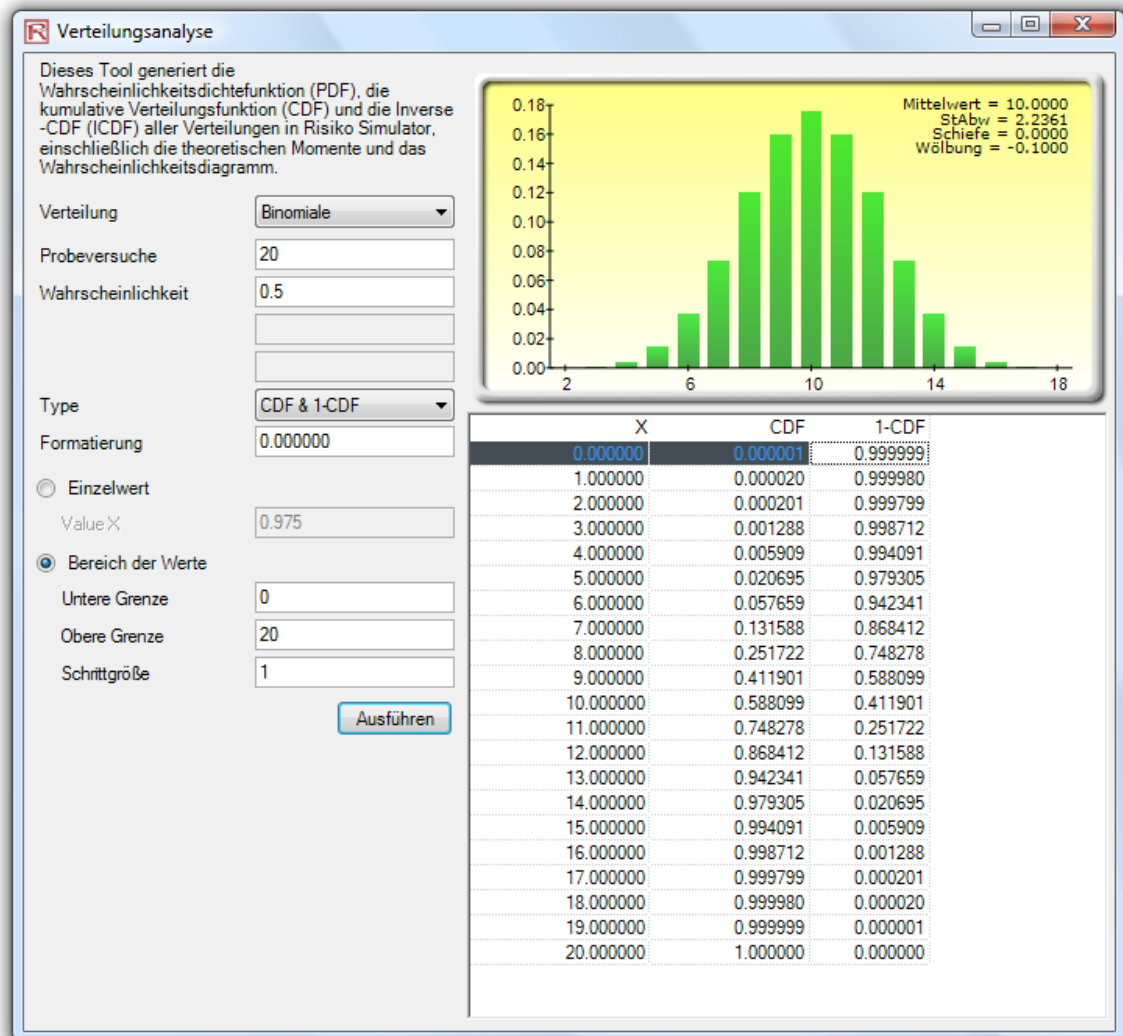


Bild 5.36—Verteilungsanalyse Tool (CDF einer Binomialverteilung mit 20 Probeversuchen)

Unter Verwendung dieses Verteilungsanalyse Tool kann man noch mehr fortgeschrittene Verteilungen analysieren, wie die Gamma-, Beta-, negative Binomial- und viele andere Verteilungen in Risiko Simulator. Als weiteres Beispiel der Verwendung des Tools in einer

kontinuierlichen Verteilung und der Funktionalität der inversen kumulativen Verteilungsfunktion (ICDF), zeigt Bild 5.37 die Standardnormalverteilung (Normalverteilung mit einem Mittelwert von Null und einer Standardabweichung von eins), wobei wir die ICDF anwenden, um den Wert von x zu finden, der die kumulative Wahrscheinlichkeit von 97.50% (CDF) entspricht. Das heißt, eine Einschwanz kumulative Verteilungsfunktion (CDF) von 97.50% ist äquivalent zu einem Zweisechwanz 95% Konfidenzintervall (es gibt eine Wahrscheinlichkeit von 2.50% im rechten Schwanz und von 2.50% im linken Schwanz, was 95% im Zentrum oder im Konfidenzintervallbereich übrig lässt, was seinerseits äquivalent zu einem 97.50% Bereich für einen Schwanz ist). Das Ergebnis ist der bekannte Z-Wert von 1.96. Deshalb kann man, unter Verwendung dieses Verteilungsanalyse Tool, sowohl die standardisierten Werte für andere Verteilungen als auch die exakten und kumulativen von anderen Verteilungen schnell und mühelos erhalten.

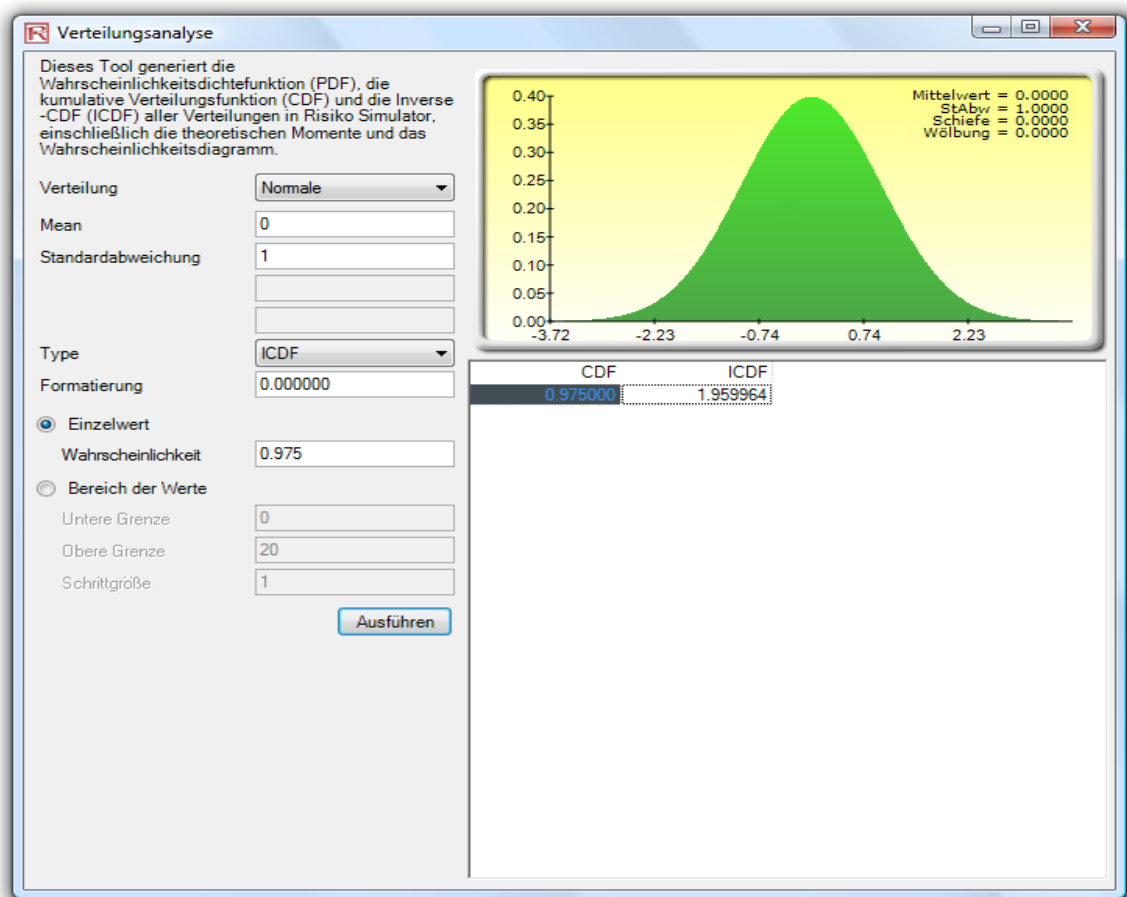


Bild 5.37—VerteilungsanalyseTool (ICDF and Z-Wert einer Normalverteilung)

Szenarioanalyse Tool

Das Szenarioanalyse Tool in Risiko Simulator erlaubt es Ihnen schnell und mühelos mehrfache Szenarien auszuführen und dabei ein oder zwei Inputparameter zu ändern, um den Output einer Variablen festzustellen. Bild 5.38 zeigt wie dieses Tool im Beispiel „diskontierter Cashflow

Modell“ funktioniert (Modell 7 im Ordner Beispielsmodelle von Risiko Simulator). In diesem Beispiel wird die Zelle G6 (Nettogegegenwartswert) als der Output der uns interessiert ausgewählt, während die Zellen C9 (effektiver Steuersatz) und C12 (Produktpreis) als die zu störenden Inputs ausgewählt werden. Sie können sowohl die zu testenden Anfangs- und Endwerte als auch die Schrittgröße oder die Anzahl der zwischen diesen Anfangs- und Endwerten auszuführenden Schritte einstellen. Das Ergebnis ist eine Szenarioanalysetabelle (Bild 5.39), wobei die Reihen- und Spaltenkopfzeilen die zwei Inputvariablen sind und der Tabellenbereich die Nettogegegenwartswerte zeigt.

A	B	C	D	E	F	G	H	I	J	K	L	M
1												
2												
3												
4	Basissjahr	2009			Summe PV Nettovorteile	\$4,762.09		Diskontyp	Diskrete End-of-Year Discounting			
5	Anfangsjahr	2009			Summe PV Investitionen	\$1,634.72						
6	Markt Risikokompakter Diskontsatz	15.00%			Nettogegegenwartswert	\$3,127.87		Modelltyp	Include Terminal Valuation			
7	Privat Risiko Diskontsatz	5.00%			Interner Zinsfuß	55.60%						
8	Endperiode Wachstumsrate	2.00%			Kapitalrendite (ROI)	191.40%						
9	Effektiver Steuersatz	40.00%			Profitabilitätsindex	2.91						
10												
11		2009	2010	2011	2012	2013	2014	2015	2016	2017	2018	
12	Produkt A Durchschnittspreis/Einheit	\$10.00	\$10.50	\$11.00	\$11.50	\$12.00	\$12.50	\$13.00	\$13.50	\$14.00	\$14.50	
13	Produkt B Durchschnittspreis/Einheit	\$12.25	\$12.50	\$12.75	\$13.00	\$13.25	\$13.50	\$13.75	\$14.00	\$14.25	\$14.50	
14	Produkt C Durchschnittspreis/Einheit	\$15.15	\$15.30	\$15.45	\$15.60							
15	Produkt A Verkaufsmenge (000e)	50	50	50								
16	Produkt B Verkaufsmenge (000e)	35	35	35								
17	Produkt C Verkaufsmenge (000e)	20	20	20								
18	Gesamteinnahmen	\$1,231.75	\$1,268.50	\$1,305.25	\$1,342.00							
19	Direktkosten der verkauften Waren	\$184.76	\$190.28	\$195.79	\$201.30							
20	Bruttogewinn	\$1,046.99	\$1,078.23	\$1,109.46	\$1,140.70							
21	Betriebsausgaben	\$157.50	\$157.50	\$157.50	\$157.50							
22	Verkaufs-, Allgemein- und Verwaltungskosten	\$15.75	\$15.75	\$15.75	\$15.75							
23	Betriebsleistungen (EBITDA)	\$873.74	\$904.98	\$936.21	\$967.44							
24	Abschreibung	\$10.00	\$10.00	\$10.00	\$10.00							
25	Amortisierung	\$3.00	\$3.00	\$3.00	\$3.00							
26	EBIT	\$860.74	\$891.98	\$923.21	\$954.44							
27	Zinszahlungen	\$2.00	\$2.00	\$2.00	\$2.00							
28	EBT	\$858.74	\$889.98	\$921.21	\$952.44							
29	Steuern	\$343.50	\$355.99	\$368.49	\$380.98							
30	Nettoeinnahmen	\$515.24	\$533.99	\$552.73	\$571.46							
31	Unbar: Abschreibung Amortisierung	\$13.00	\$13.00	\$13.00	\$13.00							
32	Unbar: Änderung im Nettobetriebskapital	\$0.00	\$0.00	\$0.00	\$0.00							
33	Unbar: Kapitalaufwendungen	\$0.00	\$0.00	\$0.00	\$0.00							
34	Freier Cashflow	\$528.24	\$546.99	\$565.73	\$584.47	\$603.21	\$621.36	\$639.50	\$657.64	\$675.78	\$5,444.64	
35												
36	Investitionsauslage	\$500.00	\$1,500.00									
37												
38	Netto freier Cashflow	(\$1,705.97)	\$546.99	\$565.73	\$584.47	\$603.21	\$621.36	\$639.50	\$657.64	\$675.78	\$5,444.64	
39												

Bild 5.38—Szenarioanalyse Tool

SZENARIOANALYSETABELLE																								
Outputvariable:	\$G\$6	Anfangsgrundfallwert:		\$3,127.87																				
Spaltenvariable:	\$C\$12	Min:	10 Max:	30 Schritte:	20 Schrittgröße:																			
Reihenvariable:	\$C\$9	Min:	0.3 Max:	0.5 Schrittgröße:	0.01	Anfangsgrundfallwert:	\$10.00																	
	\$10.00	\$11.00	\$12.00	\$13.00	\$14.00	\$15.00	\$16.00	\$17.00	\$18.00	\$19.00	\$20.00	\$21.00	\$22.00	\$23.00	\$24.00	\$25.00	\$26.00	\$27.00	\$28.00	\$29.00	\$30.00			
30.00%	\$3,904.83	\$4,134.43	\$4,364.04	\$4,593.64	\$4,823.24	\$5,052.84	\$5,282.44	\$5,512.04	\$5,741.64	\$5,971.24	\$6,200.85	\$6,430.45	\$6,660.05	\$6,889.65	\$7,119.25	\$7,348.85	\$7,578.45	\$7,808.05	\$8,037.65	\$8,267.25	\$8,496.85			
31.00%	\$3,827.14	\$4,053.46	\$4,279.78	\$4,506.10	\$4,732.42	\$4,958.74	\$5,185.06	\$5,411.39	\$5,637.71	\$5,864.03	\$6,090.35	\$6,316.67	\$6,542.99	\$6,769.31	\$6,995.63	\$7,221.95	\$7,448.28	\$7,674.60	\$7,900.92	\$8,127.24	\$8,353.56			
32.00%	\$3,749.44	\$3,972.40	\$4,195.36	\$4,418.32	\$4,641.28	\$4,864.24	\$5,087.20	\$5,310.17	\$5,533.13	\$5,756.09	\$5,979.05	\$6,202.01	\$6,425.97	\$6,648.93	\$6,872.89	\$7,095.85	\$7,318.81	\$7,541.74	\$7,764.78	\$7,987.72	\$8,210.65			
33.00%	\$3,671.75	\$3,891.51	\$4,111.27	\$4,331.03	\$4,550.79	\$4,770.55	\$4,990.31	\$5,210.07	\$5,429.83	\$5,649.59	\$5,869.35	\$6,089.11	\$6,308.87	\$6,528.63	\$6,748.39	\$6,968.15	\$7,187.91	\$7,407.67	\$7,627.43	\$7,847.21	\$8,066.97			
34.00%	\$3,594.05	\$3,810.53	\$4,027.01	\$4,243.49	\$4,459.97	\$4,676.45	\$4,892.93	\$5,109.41	\$5,325.89	\$5,542.37	\$5,758.85	\$5,975.33	\$6,191.82	\$6,408.29	\$6,624.77	\$6,841.25	\$7,057.73	\$7,274.21	\$7,490.71	\$7,707.19	\$7,923.67			
35.00%	\$3,516.35	\$3,729.55	\$3,942.76	\$4,155.96	\$4,369.16	\$4,582.36	\$4,795.56	\$5,008.76	\$5,221.96	\$5,435.16	\$5,648.36	\$5,861.57	\$6,074.77	\$6,287.97	\$6,501.17	\$6,714.37	\$6,927.57	\$7,140.77	\$7,353.97	\$7,567.17	\$7,780.38			
36.00%	\$3,438.65	\$3,648.58	\$3,858.50	\$4,068.42	\$4,278.34	\$4,488.26	\$4,698.18	\$4,908.10	\$5,118.03	\$5,327.95	\$5,537.87	\$5,747.79	\$5,957.71	\$6,167.63	\$6,377.55	\$6,587.47	\$6,797.39	\$7,007.32	\$7,217.24	\$7,427.16	\$7,637.08			
37.00%	\$3,360.95	\$3,567.60	\$3,774.24	\$3,980.88	\$4,187.53	\$4,394.17	\$4,600.81	\$4,807.45	\$5,014.09	\$5,220.73	\$5,427.37	\$5,634.01	\$5,840.65	\$6,047.29	\$6,253.94	\$6,460.58	\$6,667.22	\$6,873.86	\$7,080.50	\$7,287.14	\$7,493.78			
38.00%	\$3,283.27	\$3,486.53	\$3,689.79	\$3,893.05	\$4,096.31	\$4,299.57	\$4,502.83	\$4,706.09	\$4,910.15	\$5,113.51	\$5,316.87	\$5,520.24	\$5,723.60	\$5,926.96	\$6,130.32	\$6,333.68	\$6,537.04	\$6,740.40	\$6,943.76	\$7,147.13	\$7,350.49			
39.00%	\$3,205.57	\$3,405.65	\$3,605.73	\$3,805.81	\$4,005.89	\$4,205.97	\$4,406.05	\$4,606.14	\$4,806.22	\$5,006.30	\$5,206.38	\$5,406.46	\$5,606.54	\$5,806.62	\$6,006.70	\$6,206.78	\$6,406.87	\$6,606.95	\$6,807.03	\$7,007.11	\$7,207.19			
40.00%	\$3,127.87	\$3,324.67	\$3,521.48	\$3,718.28	\$3,915.08	\$4,111.88	\$4,308.68	\$4,505.48	\$4,702.28	\$4,899.08	\$5,095.88	\$5,292.68	\$5,489.48	\$5,686.28	\$5,883.08	\$6,079.88	\$6,276.68	\$6,473.48	\$6,670.28	\$6,867.08	\$7,063.88			
41.00%	\$3,050.18	\$3,243.70	\$3,437.22	\$3,630.74	\$3,824.26	\$4,017.78	\$4,211.30	\$4,404.82	\$4,598.35	\$4,791.87	\$4,985.39	\$5,178.91	\$5,372.43	\$5,565.95	\$5,759.47	\$5,952.99	\$6,146.51	\$6,340.03	\$6,533.55	\$6,727.08	\$6,920.60			
42.00%	\$2,972.48	\$3,162.72	\$3,352.96	\$3,543.20	\$3,733.45	\$3,923.69	\$4,113.93	\$4,304.17	\$4,494.41	\$4,684.65	\$4,874.89	\$5,065.13	\$5,255.37	\$5,445.61	\$5,635.85	\$5,826.10	\$6,016.34	\$6,206.58	\$6,396.82	\$6,587.06	\$6,777.30			
43.00%	\$2,894.79	\$3,081.75	\$3,268.71	\$3,455.67	\$3,642.63	\$3,829.59	\$4,016.55	\$4,203.51	\$4,390.47	\$4,577.43	\$4,764.40	\$4,951.36	\$5,138.32	\$5,325.28	\$5,512.24	\$5,699.20	\$5,886.16	\$6,073.12	\$6,260.08	\$6,447.04	\$6,634.00			
44.00%	\$2,817.09	\$3,000.77	\$3,184.45	\$3,368.13	\$3,551.81	\$3,735.49	\$3,919.18	\$4,102.86	\$4,286.54	\$4,470.22	\$4,653.90	\$4,837.58	\$5,021.26	\$5,204.94	\$5,388.62	\$5,572.30	\$5,756.98	\$5,941.66	\$6,126.34	\$6,311.02	\$6,495.70			
45.00%	\$2,739.39	\$2,919.79	\$3,100.20	\$3,280.60	\$3,461.00	\$3,641.40	\$3,821.80	\$4,002.20	\$4,182.60	\$4,363.00	\$4,543.40	\$4,723.80	\$4,904.20	\$5,084.61	\$5,265.01	\$5,445.41	\$5,625.81	\$5,806.21	\$5,986.61	\$6,167.01	\$6,347.41			
46.00%	\$2,661.70	\$2,838.82	\$3,015.94	\$3,193.06	\$3,370.18	\$3,547.30	\$3,724.42	\$3,901.54	\$4,078.66	\$4,255.78	\$4,432.91	\$4,610.03	\$4,787.15	\$4,964.27	\$5,141.39	\$5,318.51	\$5,495.63	\$5,672.75	\$5,849.87	\$6,026.99	\$6,204.12			
47.00%	\$2,584.00	\$2,757.84	\$2,931.68	\$3,105.52	\$3,279.37	\$3,453.21	\$3,627.05	\$3,800.89	\$3,974.73	\$4,148.57	\$4,322.41	\$4,496.25	\$4,670.09	\$4,843.93	\$5,017.77	\$5,191.62	\$5,365.46	\$5,539.30	\$5,713.14	\$5,886.98	\$6,060.82			
48.00%	\$2,506.31	\$2,676.97	\$2,847.63	\$3,017.99	\$3,188.55	\$3,359.11	\$3,529.67	\$3,700.23	\$3,870.79	\$4,041.35	\$4,211.91	\$4,382.48	\$4,553.04	\$4,723.60	\$4,894.16	\$5,064.72	\$5,235.28	\$5,405.84	\$5,576.40	\$5,746.96	\$5,917.52			
49.00%	\$2,428.61	\$2,595.89	\$2,763.17	\$2,930.45	\$3,097.73	\$3,265.01	\$3,432.29	\$3,599.58	\$3,766.86	\$3,934.14	\$4,101.42	\$4,268.70	\$4,435.98	\$4,603.26	\$4,770.54	\$4,937.82	\$5,105.10	\$5,272.38	\$5,439.66	\$5,606.94	\$5,774.22			
50.00%	\$2,350.91	\$2,514.91	\$2,678.92	\$2,842.92	\$3,006.92	\$3,170.92	\$3,334.92	\$3,498.92	\$3,662.92	\$3,826.92	\$3,990.92	\$4,154.92	\$4,318.92	\$4,482.92	\$4,646.93	\$4,810.93	\$4,974.93	\$5,138.93	\$5,302.93	\$5,466.93	\$5,630.93			

Bild 5.39—Szenarioanalysetabelle

Segmentierung Clustering Tool

Ein letztes analytisches Verfahren von Interesse ist das der Segmentierung - Clustering. Bild 6.25 zeigt einen Beispieldatensatz. Sie können die Daten auswählen und das Tool unter **Risiko Simulator | Tools | Segmentierung Clustering** ausführen. Im Bild 5.40 zeigen wir das Beispiel einer Segmentierung von zwei Gruppen. Das heißt, wir nehmen den Originaldatensatz und führen einige interne Algorithmen aus (eine Kombination oder k-Means hierarchische Clustering und andere Methoden von Momenten, um die bestpassenden Gruppen oder natürlichen statistischen Clusters zu finden), um den Originaldatensatz statistisch in zwei Gruppen zu teilen oder zu segmentieren. Sie können die Zwei-Gruppen Mitgliedschaften im Bild 5.40 sehen. Natürlich können Sie diesen Datensatz in beliebig viele Gruppen segmentieren. Dieses Verfahren ist nützlich in einer Vielfalt von Bereichen, einschließlich Marketing (Marktsegmentierung von Kunden in verschiedenen Kundenverhältnis-Managementgruppen und so weiter), Naturwissenschaft, Ingenieurwissenschaft und andere.

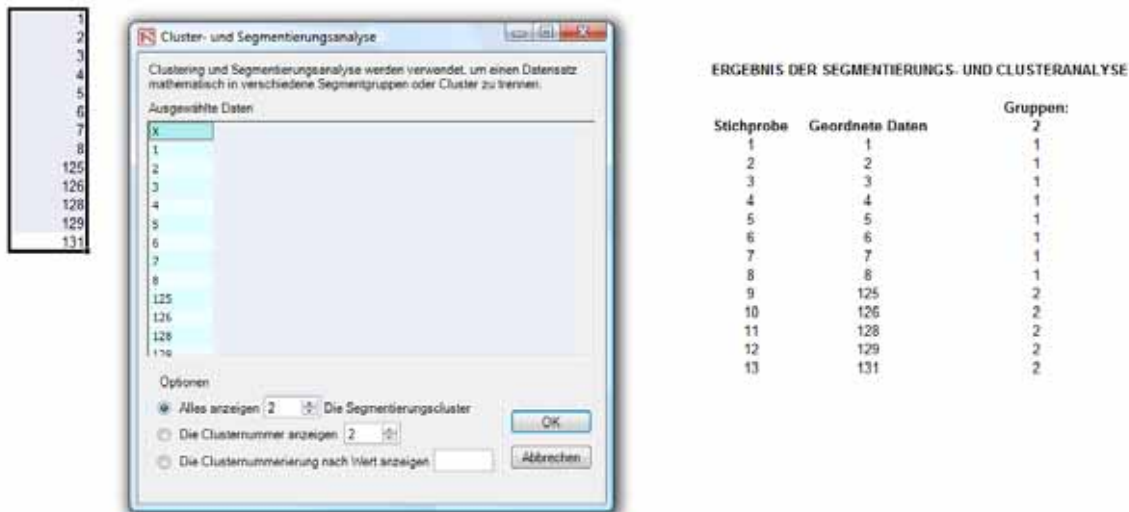


Bild 5.40—Segmentierung Clustering Tool und Ergebnisse

RISIKO SIMULATOR 2011/2012 NEUE WERKZEUGE

Zufallszahl Generierung, Monte Carlo versus Latin Hypercube, und Korrelation Kopula Methoden

Angefangen mit Version 2011/2012, gibt es 6 Zufallszahl-Generatoren, 3 Korrelation Kopulae, und 2 Simulations Sampling Methoden zur Auswahl (Bild 5.41). Diese Präferenzen sind durch die *Risiko Simulator | Optionen* Speicherstelle eingestellt einzustellen

Die Zufallszahl-Generierung (RNG) findet sich im Herzen jeder Simulations-Software. Basierend auf der generierten Zufallszahl, können verschiedene mathematische Verteilungen konstruiert werden. Das Standardverfahren ist die ROV Risiko Simulator proprietäre Methodik, welche die besten und robustesten Zufallszahlen anbietet. Es gibt 6 unterstützte Zufallszahl Generatoren, und im Allgemeinen, das ROV Risiko Simulator Standardverfahren und die Advanced Subtractive Random Shuffle Methodik (erweiterte subtraktive Zufallsmischung Methodik) sind die zwei empfohlenen Herangehensweisen. Die andere Methoden sind nicht zu verwenden, außer wenn Ihr Model oder Analytik dies spezifisch verlangt, und auch wenn, empfehlen wir die Ergebnisse gegen diese zwei vorgeschlagenen Herangehensweisen zu testen. Um so weiter unten in der Liste der RNGs, um so einfacher der Algorithmus und um so schneller läuft er, gegenüber dem weiter oben in der Liste der RNGs sind die Ergebnisse robuster

Im Korrelationsbereich sind drei Methoden unterstützt: die Normale Kopula, T-Kopula und Quasi-Normal Kopula. Diese Methoden sind auf mathematischen Integrationstechniken angewiesen, und im Zweifel bietet die normale Kopula die sichersten und konservativsten Ergebnisse. Die T-Kopula gewährleistet extreme Werte am Ende der simulierten Verteilungen, wobei die Quasi-Normal Kopula bringt Ergebnisse die zwischen diese Werte.

Im Bereich der Simulationsmethodik, werden Monte Carlo Simulation (MCS) und Latin Hypercube Sampling (LHS) Methoden unterstützt. Beachten Sie, dass Kopulas und andere multivariate Funktionen sind **nicht** mit LHS kompatibel. Der Grund hierfür ist, dass LHS bei einer einzelnen Zufallsvariable verwendet werden kann, aber nicht auf einer gemeinsamen Verteilung. In Wirklichkeit, hat LHS sehr begrenzte Beeinflussung auf die Genauigkeit der Ausgabe je mehr Verteilungen in einem Model sind, weil LHS nur anwendbar bei individuellen Verteilungen ist. Der Nutzen des LHS ist ebenso erodiert wenn man die Zahl der Samples am Anfang nicht vervollständigt, z.B. wenn man die Simulationslauf mittendrin anhält. LHS übernimmt auch eine schwere Belastung auf einem Simulationsmodel mit einer großen Eingabezahl, weil es notwendig ist, Samples aus jeder Verteilung zu generieren und organisieren bevor das erste Sample aus einer Verteilung abläuft. Dies kann eine lange Verzögerung in dem Ablaufen eines großen Models verursachen, bietet jedoch sehr wenig zusätzliche Genauigkeit. Letztendlich ist LHS am besten anzuwenden wenn die Verteilungen artig, symmetrisch und ohne jegliche Korrelationen sind. Trotzdem ist LHS eine leistungsstarke Herangehensweise die eine

gleichmäßig abgetastete Verteilung ergibt, wo MCS manchmal stückige Verteilungen generieren kann (abgetastete Werte können manchmal stärker in einem Verteilungsbereich konzentriert sein) im Vergleich zu einer gleichmäßig abgetasteten Verteilung (jedes Teil der Verteilung wird abgetastet) wenn LHS angewandt wird.

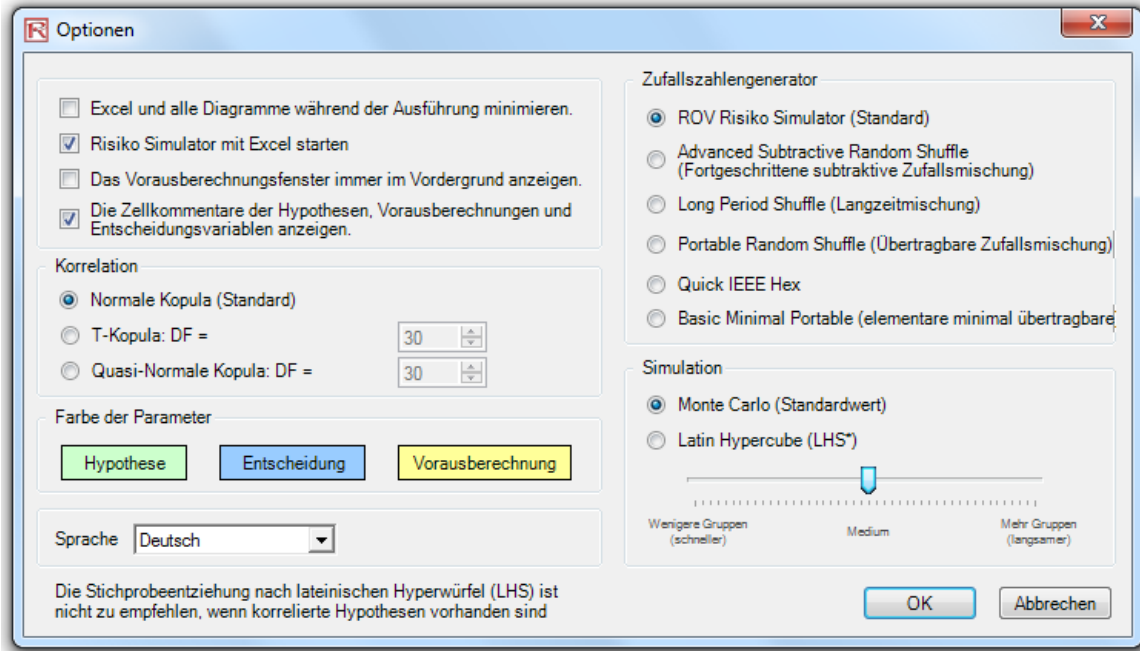


Bild 5.41—Risiko Simulator Optionen

Daten Saisonbereinigung und Trendbereinigung

Dieses Tool bereinigt saisonal und säubert Ihre Grunddaten von Trends, um jegliche saisonale und trending Komponente zu entfernen (Bild 5.42). Normalerweise beinhaltet der Prozess in Prognosemodellen die Entfernung der Auswirkungen von akkumulierenden Datensätze von Saisonalität und Trend, um nur die absoluten Wertänderungen zu zeigen, und potentielle Konjunkturschemen zyklische Muster zu gewähren nachdem die allgemeine Abweichungen, Tendenzen, Wendungen, Biegungen und Auswirkungen der saisonalen Zyklen einem Zeitreihendatensatz entfernt sind Zum Beispiel, ein bereinigter Datensatz könnte notwendig sein um eine genauere Berechnung der Betriebsumsätze eines vorgegebenen Jahres zu sehen, wenn der ganzen Datensatz von einer Neigung zu einer flachen Ebene verschoben wird um die grundlegenden Zyklen und Schwankungen zu erkennen.

Viele Zeitreihendaten weisen Saisonalität auf wo sich bestimmte Vorgänge wiederholen nach einem Zeitraum oder Saisonalität Periode (z.B. Skiort Umsätze sind im Winter höher als im Sommer, und dieser berechenbare Zyklus wird sich jeden Winter wiederholen.) Saisonalität Perioden representieren wie viele Perioden vorbeigehen müssten bevor der Zyklus sich wiederholt (z.B. 24 Stunden in einem Tag, 12 Monate in einem Jahr, 4 Quartale in einem Jahr, 60 Minuten in

einer Stunde, und so weiter). Dieses Tool bereinigt saisonal und entfernt Trends in Ihren Grunddaten um jeglichen saisonalen Komponenten zu eliminieren. Ein saisonaler Index größer als 1 zeigt einen Höhepunkt innerhalb des Zyklus, und ein Wert unter 1 zeigt einen Abfall im Zyklus.

Ablauf (Daten Saisonbereinigung und Trendbereinigung):

- Zu analysierende Daten auswählen (z.B. B9:B28) und klicken Sie auf **Risiko Simulator | Tools | Daten Saisonbereinigung und Trendbereinigung**
- **Daten Saisonbereinigung** und/oder **Trendbereinigung** auswählen, die Trendbereinigungsverfahren die Sie laufen lassen möchten aussuchen, in die relevanten Order eingeben (z.B. polynomial Order, gleitender Durchschnitt Order, Differenz Order und Rate Order, dann OK klicken
- Die zwei generierte Berichte durchsehen für weitere Details über Methodologie, Anwendung, und resultierende Diagramme und saisonalbereinigt/trendbereinigte Daten.

Ablauf (Saisonbereinigung und Trendbereinigung):

- Zu analysierende Daten auswählen (z.B. B9:B28) und klicken Sie auf **Risiko Simulator | Tools | Daten Saisonalität Test**.
- In die höchste Saisonalitäts-Periode eingeben um auszutesten. Das heißt, wenn Sie 6 eingeben, wird das Tool folgende Saisonalitäts-Perioden testen: 1,2,3,4,5,6. Periode 1 impliziert natürlich keine Saisonalität in den Daten.
- Den generierten Bericht für weitere Details über Methodologie, Anwendung, resultierende Diagramme und Saisonalitäts-Testergebnisse durchsehen. Die beste Saisonalitäts Periodizität ist zuerst aufgelistet (rangiert nach der niedrigsten RMSE Fehlermaß) und alle relevante Fehlermessungen sind einbezogen zu Vergleichszwecken: root mean squared error (RMSE)—quadratischer Mittelwert Fehler, mean squared error (MSE)—mittlerer quadratischer Fehler (MQF), mean absolute deviation (MAD)—durchschnittliche absolute Abweichung, und mean absolute percentage error (MAPE) —durchschnittliche absolute prozentuelle Fehler.

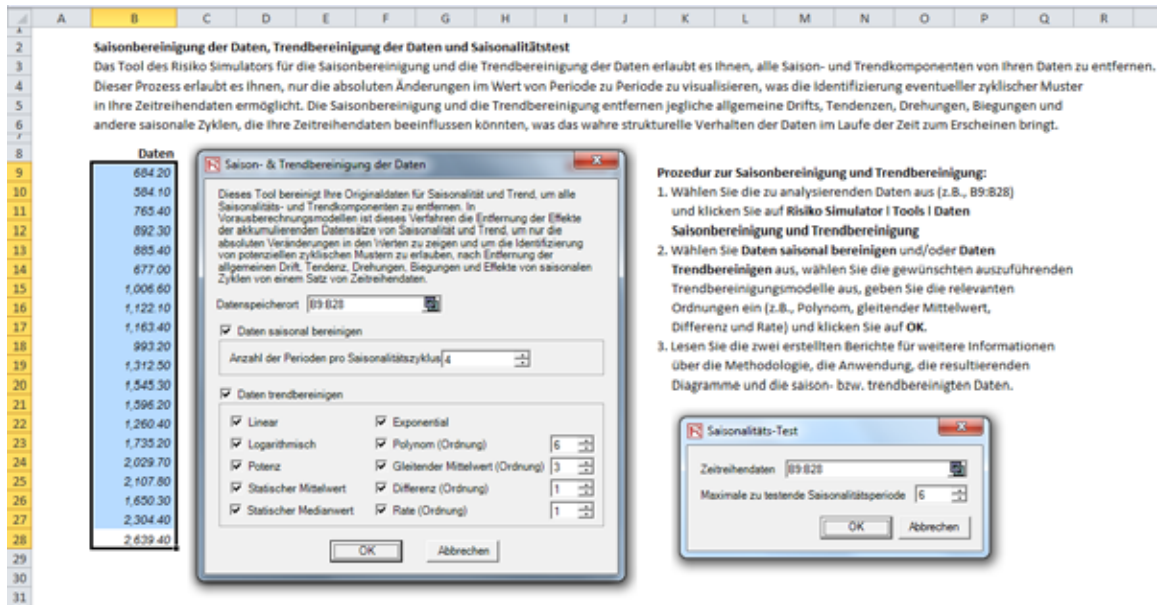


Bild 5.42 - Daten Saisonalbereinigung und Trendbereinigung

Hauptkomponentenanalyse

Die Hauptkomponentenanalyse ist eine Methode um Datenschemata zu identifizieren und die Daten umzugestalten, so dass die Gemeinsamkeiten und Verschiedenheiten gehighlighted hervorgehoben sind (Bild 5.43). Datenschemata sind sehr schwer zu finden wenn multiple mehrere Variablen existieren, und höherdimensionale Graphen sind sehr schwierig darzustellen und interpretieren. Sobald die Schemata gefunden sind, können sie komprimiert werden, und die Anzahl der Dimensionen ist jetzt reduziert. Diese Reduktion der Datendimensionen bedeutet nicht viel Reduktion des Informationsverlusts. Stattdessen, ähnliche Informationsebenen können jetzt mit niedriger Variablenanzahl erzielt werden.

Ablauf: Procedure:

- Zu analysierende Daten auswählen (z.B. B11:K30), klicken Sie auf **Risiko Simulator | Tools | Hauptkomponentenanalyse** und
- Den generierten Bericht nach den errechneten Ergebnissen überprüfen

VAR1	VAR2	VAR3	VAR4	VAR5	VAR6	VAR7	VAR8	VAR9	VAR10	Prozedur:
96.998	87.223	102.443	112.765	111.984	117.331	78.164	97.658	110.950	89.133	1. Wählen Sie die zu analysierenden Daten aus (z.B., B11:K30), klicken Sie auf
93.098	83.096	81.531	90.224	92.265	78.821	94.321	95.960	101.349	96.345	Risiko Simulator Tools Hauptkomponenten-Analyse
96.730	96.298	113.426	99.147	98.138	94.868	119.722	108.657	123.757	93.451	und klicken Sie auf OK .
116.615	83.876	105.389	109.022	119.189	99.155	94.762	106.751	96.187	107.576	2. Lesen Sie den erstellten Bericht für die errechneten Ergebnisse.
85.558	91.528	84.784	96.371	99.675	100.281	96.773	121.945	82.575	92.635	
74.224	114.477	87.202	93.464	107.577	104.667	108.746	105.957	86.282	88.843	
106.940	103.226	90.602	97.591	101.315	105.578	101.387	90.890	118.848	104.872	
100.722	108.298	108.620	93.635	90.768	111.112	87.988	84.411	107.113	106.384	
122.057	114.438	113.039	101.130	100.020	104.537	99.745	89.453	82.252	108.283	
104.442	106.179	102.135	89.731	112.382	96.888	91.601	91.789	95.710	95.466	
94.762	108.494	105.132	93.917	113.050	82.391	105.506	98.837	100.417	93.459	
94.504	108.493	108.030	104.564	106.914	116.306	103.039	105.890	118.528	96.644	
110.383	101.435	111.410	98.517	92.202	110.760	94.182	105.339	105.458	96.836	
95.592	86.340	119.930	94.335	100.861	97.657	128.354	112.520	108.809	113.322	
101.879	105.420	97.504	87.789	112.667	97.111	86.941	107.643	107.843	104.282	
104.039	93.519	107.231	105.253	110.750	72.306	104.638	114.671	82.774	100.455	
113.540	116.882	102.387	101.451	118.545	99.574	93.431	109.074	99.901	110.392	
104.347	114.534	98.788	90.383	84.614	74.349	101.032	102.992	99.822	102.005	
102.582	114.762	100.853	88.833	86.101	101.915	109.511	84.912	93.900	105.235	
97.832	96.564	98.365	95.603	91.974	106.448	100.588	112.635	102.622	100.571	

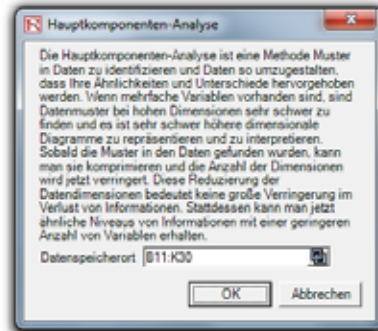


Bild 5.43 - Hauptkomponentenanalyse

Strukturbruchanalyse

Ein Strukturbruch prüft ob die Koeffizienten in verschiedenen Datensätzen gleichwertig sind, und dieser Test ist am häufigsten in der Zeitreihenanalyse verwendet um die Gegenwart einem Strukturbruch zu prüfen (Bild 5.44). Ein Zeitreihendatensatz kann in zwei Teilsätze aufgeteilt werden, und jeder Teilsatz wird auf dem anderen, sowie auf dem gesamten Datensatz, getestet um statistisch festzustellen ob tatsächlich ein Bruch zu einer bestimmten Zeitspanne beginnt. Der Strukturbruchtest wird oft angewendet um festzustellen ob die unabhängige Variablen verschiedene Einflüsse auf verschiedene Untergruppen der Population haben, wie zum Beispiel um auszutesten ob eine neue Marketingkampagne, Aktivität, Großveranstaltung, Akquisition, Veräußerung, und so weiter, auf den Zeitreihendaten Einfluß nimmt. Angenommen der Datensatz hat 100 Zeitreihen Datenpunkte, können Sie verschiedene Stopppunkte einrichten um Datenpunkte 10, 30 und 51 zum Beispiel, zu testen (dies bedeutet, dass drei strukturierte Abbruchtests auf dem folgenden Datensatz aufgeführt werden: Datenpunkte 1-9 verglichen mit 10-100; Datenpunkte 1-29 verglichen mit 30-100; und 1-50 verglichen mit 51-100, um zu sehen ob sich tatsächlich am Anfang des Datenpunktes 10, 30 und 51 ein Abbruch in der unterliegenden Struktur befindet). Ein einseitiger Hypothesentest wird auf der null Hypothese (H_0) aufgeführt, insofern als die zwei Datenteilsätze statistisch miteinander vergleichbar sind. Das heißt, es gibt keinen statistisch signifikanten Strukturbruch. Die alternativ Hypothese (H_a) ist, dass die zwei Datenteilsätze statistisch unterschiedlich voneinander sind, was auf einen möglichen Strukturbruch hindeutet. Wenn die kalkulierte p-Werte weniger als oder gleichwertig mit 0.01, 0.05, oder 0.10 sind, bedeutet dies, dass die Hypothese verworfen wird, was impliziert, dass die zwei Datenteilsätze statistisch bezeichnenderweise an die 1%, 5% und 10% Signifikanzniveaus verschieden sind. Hohe p-Werte bedeutet, dass es keinen statistisch signifikanten Strukturbruch gibt.

Ablauf Procedure:

- Zu analysierende Daten auswählen (z.B. B15:D34), klicken Sie auf **Risiko Simulator | Tools | Strukturbruch Test** und eingeben in die relevanten Testpunkte die Sie auf die Daten anwenden möchten (z.B. 6, 10, 12) und OK klicken.
- Den generierten Bericht überprüfen um festzustellen welche dieser Testpunkte auf einen statistisch signifikanten Bruchpunkt in Ihren Daten hindeuten, und welche Punkte nicht.

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

Prozedur:

1. Wählen Sie die zu analysierenden Daten aus (z.B., B15:D34), klicken Sie auf **Risiko Simulator | Tools | Strukturbruch-Test**, geben Sie die relevanten Testpunkte ein, die Sie auf den Daten anwenden möchten (z.B., 6, 10, 12) und klicken Sie auf **OK**.
2. Lesen Sie den Bericht, um festzustellen, welcher dieser Testpunkte auf einen statistisch signifikanten Bruchpunkt in Ihren Daten deutet und welche es nicht tun.

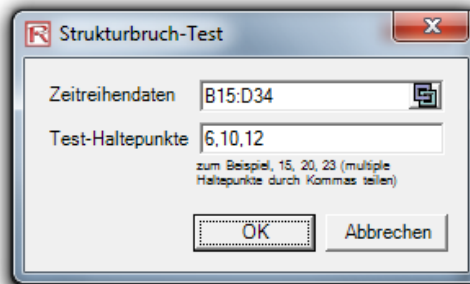


Bild 5.44—Strukturbruchanalyse

Trendlinie Voraussagen

Trendlinien können eingesetzt werden um festzustellen ob ein Zeitreihendatensatz irgendeinen abschätzbaren Trend befolgt (Bild 5.45). Trends können linear oder nichtlinear (wie expotentiell, logarithmisch, gleitender Durchschnitt, Power, polynomial, oder Power).

Ablauf: Procedure:

- Zu analysierende Daten auswählen und klicken Sie auf **Risiko Simulator | Voraussagen | Trendlinie**, die relevante Trendlinien die Sie auf die Daten anwenden möchten auswählen (z.B. alle Methoden standardmäßig auswählen), eingeben in die Zahl der zu voraussagenden Perioden (z.B. 6 Perioden) und OK klicken.
- Den generierten Bericht überprüfen um festzustellen welche diese Testtrendlinien für Ihre Daten die beste Anpassung und beste Prognose bietet.

Historische Verkaufseinnahmen

Jahr	Vierteljahr	Periode	Umsatz
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

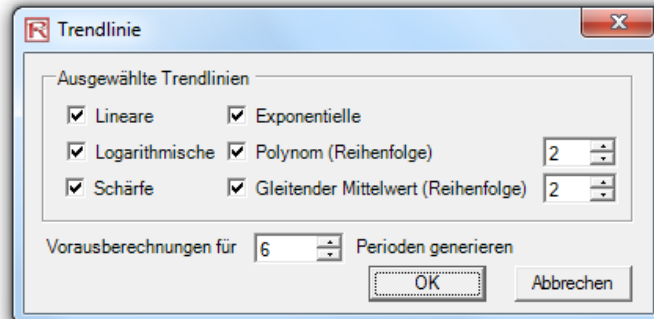


Bild 5.45—Trendlinie Voraussagen

Modellprüfungstool

Nachdem ein Modell erstellt ist, und nachdem Annahmen und Prognosen eingestellt sind, können Sie die Simulation wie gewöhnlich laufen lassen, oder das Modellprüfungstool ausführen (Bild 5.46) um zu testen ob das Modell richtig eingestellt wurde. Wenn das Modell andererseits nicht läuft, und Sie vermuten, dass einige Einstellungen falsch sind, lassen Sie dieses Tool von **Risiko Simulator | Tools | Modellprüfung** laufen um zu ermitteln wo die Probleme mit Ihrem Modell sich befinden könnten. Beachten Sie, dass dieses Tool auf die häufigste Modellprobleme sowie auf Probleme in Risiko Simulator Annahmen und Prognosen überprüft, und ist in keiner Weise umfangreich genug um auf alle Problemformen zu überprüfen.... es liegt immer noch an dem Modellentwickler um zu versichern, dass das Modell einwandfrei funktioniert.

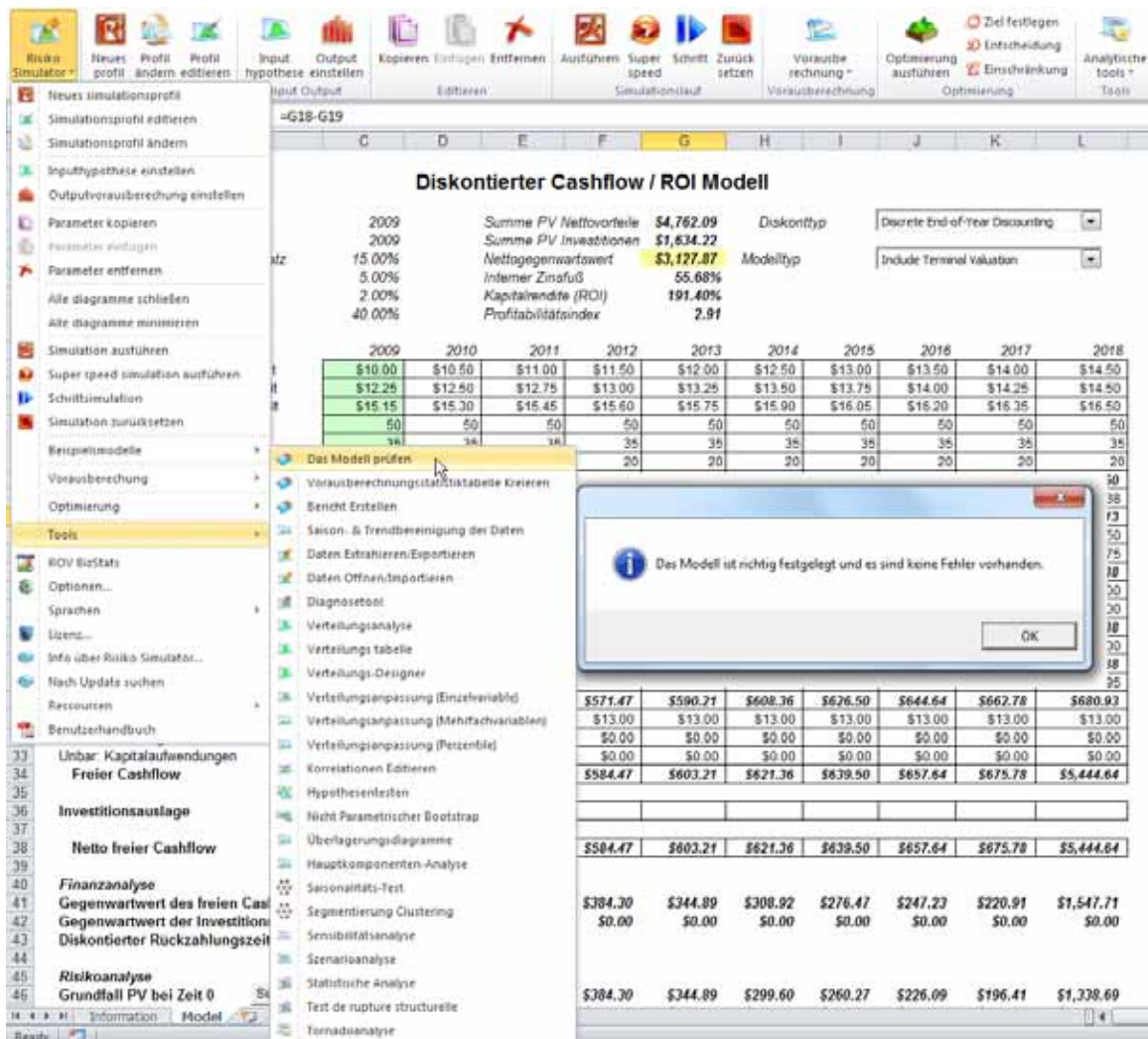


Bild 5.46—Modellprüfungstool

Perzentiles Verteilungsanpassungs-Tool

Das perzentiles Verteilungsanpassungs-Tool (Bild 5.47) ist eine weitere Methode der Wahrscheinlichkeits- Verteilungsanpassung. Es gibt einige verwandte Tools, und jede hat seine eigene Verwendungen und Vorteile:

- Verteilungsanpassung (Perzentile)—mittels alternative Eingabemethoden (Perzentile und erste/zweite Momentkombinationen) können Sie die am besten geeigneten Parameter einer bestimmten Verteilung finden, ohne die Notwendigkeit unverarbeitete Daten zu haben. Diese Methode is geeignet zum Gebrauch wenn unzureichende Daten vorhanden sind, nur wenn Perzentile und Momente verfügbar sind, oder als ein Mittel um die gesamte Verteilung wiederherzustellen mit nur zwei oder three Datapunkte aber der Verteilungstyp muss angenommen werden oder bekannt sein.
- Verteilungsanpassung (Einzelvariable)—Nutzung statistischer Methoden um die Rohdaten an alle 42 Verteilungen anzupassen, um die am besten geeignete Verteilung und

deren Eingabeparameter zu finden. Multiple Datenpunkte sind für eine gute Anpassung erforderlich, und die Verteilungstyp könnte oder könnte nicht vorzeitig bekannt sein.

- Verteilungsanpassung (Multiple Variablen)—Nutzung statistischer Methoden um Ihre Rohdaten gleichzeitig an multiple Variablen anzupassen, Nutzung der gleichen Algorithmen als die einzeln Variable-Anpassung, aber integriert eine paarweise Korrelationsmatrix zwischen den Variablen. Multiple Datenpunkte sind für eine gute Anpassung erforderlich, und die Verteilungstyp könnte oder könnte nicht vorzeitig bekannt sein.
- Benutzerdefinierte Verteilung (Festgesetzte Annahme)—Nutzung nicht parametrischer Resamplingtechniken um eine benutzerdefinierte Verteilung zu generieren mit den existierenden Rohdaten, und die Verteilung zu simulieren basierend auf dieser empirischen Verteilung. Weniger Datenpunkte sind erforderlich und die Verteilungstyp nicht vorzeitig bekannt ist.

Ablauf:

Risiko Simulator | Tools | Verteilungsanpassung (Perzentile) anklicken, Wahrscheinlichkeitsverteilung und Eingabetypen die Sie anwenden möchten auswählen, die Parameter eingeben und **Ausführen** klicken um die Ergebnisse zu bekommen. Die angepasste R-Quadrat Ergebnisse überprüfen und die empirische gegen die theoretische Anpassungsergebnisse vergleichen, um festzustellen ob Ihre Verteilung gut zusammenpasst.

Data Fitting - Subject Matter Expert Curve Fit

Diese Datenanpassungsmethode erlaubt es Ihnen, benutzerdefinierte Perzentile an Stelle einer oder mehreren normalen Inputparameter einzugeben, um die theoretische Verteilung zu bestimmen. Sie ist nützlich, wenn Expertenmeinungen über die Materie angefordert werden. Zum Beispiel, anstatt den Mittelwert und die Standardabweichung für eine Normalverteilung einzugeben, können Sie einen oder beide dieser Parameter durch Ihre eigenen Perzentile ersetzen und dieses Tool wird eine Anpassung durchführen, um die entsprechenden Verteilungsparameter zu erhalten. Kontinuierliche Verteilungen neigen dazu, sich besser als diskrete Verteilungen anzupassen (normalerweise gibt es multiple diskrete Verteilungskonfigurationen, welche sich demselben Satz von Inputs anpassen können).

Schritt 1: Wählen Sie den Verteilungs- und den Parameterbewertungstyp aus

- Dreieckige -- Minimum, Perzentil, Maximum
- Dreieckige -- Minimum, Höchstwahrscheinlich, Perzentil
- Dreieckige -- Perzentil, Perzentil, Maximum
- Dreieckige -- Perzentil, Höchstwahrscheinlich, Perzentil
- Dreieckige -- Minimum, Perzentil, Perzentil
- Dreieckige -- Perzentil, Perzentil, Perzentil
- Dreieckige -- Mittelwert, St-Abw, Perzentil
- * Uniforme *
- Uniforme -- Minimum, Perzentil
- Uniforme -- Perzentil, Maximum
- Uniforme -- Perzentil, Perzentil
- Uniforme -- Mittelwert, St-Abw
- * Weibull *
- Weibull -- Alpha, Perzentil
- Weibull -- Perzentil, Beta
- Weibull -- Perzentil, Perzentil
- Weibull -- Mittelwert, St-Abw
- * Weibull 3 *
- Weibull 3 -- Perzentil, Beta, Lage
- Weibull 3 -- Alpha, Perzentil, Lage
- Weibull 3 -- Alpha, Beta, Perzentil
- Weibull 3 -- Perzentil, Perzentil, Lage
- Weibull 3 -- Perzentil, Beta, Perzentil
- Weibull 3 -- Alpha, Perzentil, Perzentil
- Weibull 3 -- Perzentil, Perzentil, Perzentil
- Weibull 3 -- Mittelwert, St-Abw, Perzentil

Schritt 2: Geben Sie die relevanten Inputs ein

Parameter	Wert	Perzentil (%)
Perzentil	2.53	10
Perzentil	2.66	45
Perzentil	3.89	95

Schritt 3: Führen Sie die Kurvenanpassung aus und vergleichen Sie die empirische und die theoretische Verteilungen

Angepasstes R-Quadrat	
Alpha	100.0000%
Beta	0.7113
Lage	0.2935

Empirisch (Benutzer-Input)	Theoretisch (angepasst)
2.5300	2.5300
2.6600	2.6600
3.8900	3.8900
Mittelwert	2.8836
St-Abw	0.5255
Schiefte	3.4054
Wölbung	19.3606

Sprache: Deutsch Dezimale: 4

Ausführen **Schließen**

Bild 5.47—Perzentiles Verteilungsanpassungs-Tool

Verteilungscharts und Tabellen: Wahrscheinlichkeitsverteilungs Tool

Diese neue Wahrscheinlichkeitsverteilungs Tool ist ein sehr starkes und schnelles Modul angewendet für die Generierung von Verteilungscharts und Tabellen (Bild 5.48-5.51). Beachten Sie, dass sich drei ähnliche Tools im Risiko Simulator befinden, aber jedes tut ganz andere Dinge:

- Verteilungsanalyse—angewendet um schnell die PDF, CDF und ICDF von den 42 Wahrscheinlichkeitsverteilungen verfügbar im Risiko Simulator zu berechnen, und eine Wahrscheinlichkeitstabelle dieser Werte zu wiedergeben.
- Verteilungscharts und Tabellen—dieses Wahrscheinlichkeitsverteilungs Tool wird hier beschrieben und wird verwendet um verschiedene Parameter von der gleichen Verteilung zu vergleichen (z.B. Die Formen und PDF, CDF, ICDF Werte einer Weibull Verteilung mit Alpha und Beta von [2, 2], [3, 5], and [3.5, 8], und blendet sie übereinander.)
- Overlay Charts - verwendet um verschiedene Verteilungen zu vergleichen (theoretische Input-Annahmen und empirisch simulierte Output-Voraussagen), und sie aufeinanderlegen um eine visuelle Vergleichung zu erzielen.

Ablauf: Procedure:

ROV BizStats über *Risiko Simulator | Verteilungscharts und Tabellen* laufen lassen, auf dem *Globale Inputs Anwenden* Schaltfläche klicken um einen Mustersatz der Input Parameter zu laden, oder Ihre eigene Inputs eingeben und *Ausführen* klicken um die Ergebnisse zu errechnen. Die resultierende vier Momente und CDF, ICDF, PDF sind für jede der 45 Wahrscheinlichkeitsverteilungen errechnet (Bild 5.48).

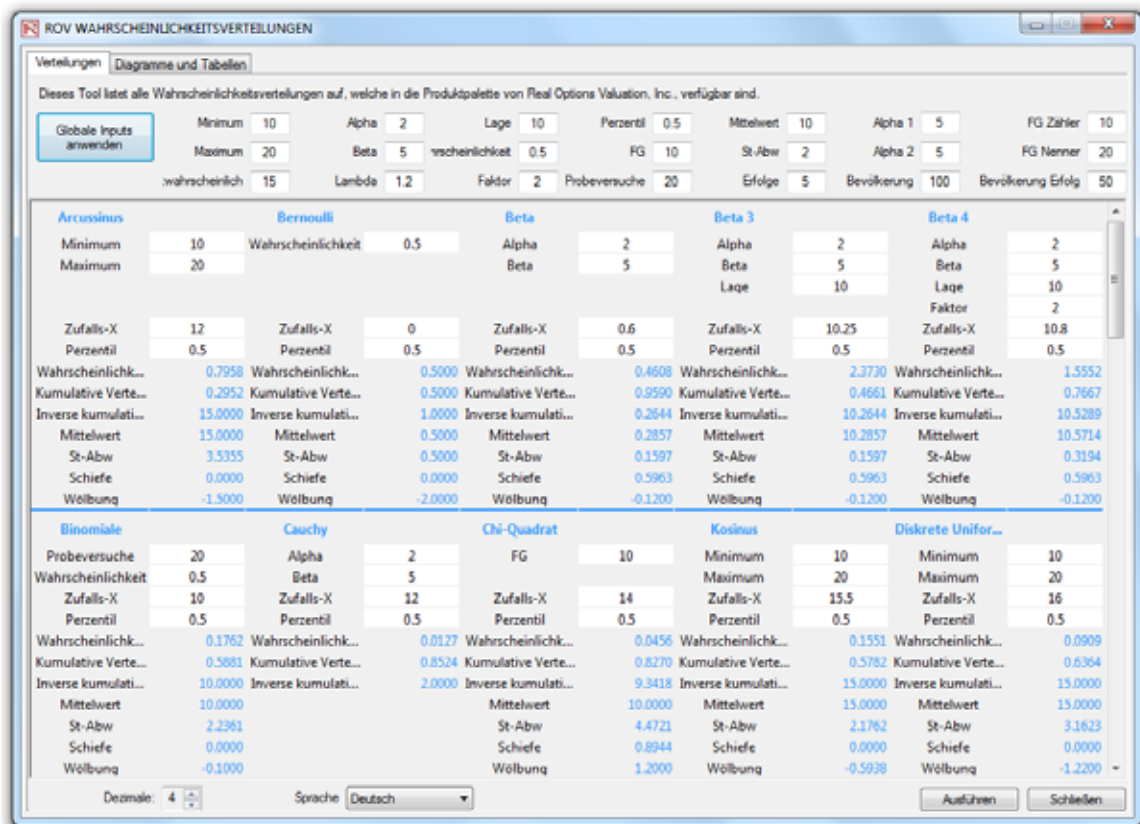


Bild 5.48 Wahrscheinlichkeitsverteilungs Tool (45 Wahrscheinlichkeitsverteilungen)

- Auf **Charts and Tabellen** Tab (Bild 5.49) klicken, eine Verteilung auswählen [A] (z.B. Arkussinus), auswählen ob Sie CDF, ICDF oder PDF laufen lassen möchten [B], relevante Inputs eingeben und **Chart Ausführen** oder **Tabelle Ausführen** [C] anklicken. Sie können zwischen dem Chart und Tabellen Tab wechseln um die Ergebnisse anzusehen, sowie einige Chart Ikonen ausprobieren [E] um die Auswirkungen auf dem Chart zu sehen.
- Sie können auch zwei Parameter ändern [H] um multipel Charts und Verteilungstabellen zu generieren wenn Sie den **Von/Zu/Stufe** Input eingeben, oder die **Benutzerdefiniert** Inputs verwenden und **Ausführen** drücken. Zum Beispiel, wie in Bild 5.50 dargestellt, die Beta Verteilung laufen lassen und PDF auswählen [G], Alpha und Beta auswählen um zu ändern [H] unter Verwendung von Custom **Benutzerdefiniert** [I] Inputs, und die relevante Input Parameter eingeben: 2;5;5 für Alpha und 5;3;5 für Beta [J] dann **Chart Ausführen** eingeben. Dieses wird drei Beta Verteilungen generieren [K]: Beta (2,5), Beta (5,3) und Beta (5,5) [L]. Verschiedene Charttypen, Rasterlinien, Sprache und Dezimaleinstellungen untersuchen, und versuchen die Verteilung nochmal laufen zu lassen unter Verwendung von theoretisch gegen empirisch simulierte Werte [N].

- Bild 5.51 veranschaulicht die Wahrscheinlichkeitstabellen generiert für eine binomiale Verteilung wo die Wahrscheinlichkeit des Erfolgs und Zahl der erfolgreichen Probeläufe (Zufallsvariable X) ausgewählt sind um zu variieren [O] unter Anwendung der *Von/Zu/Stufe* Option. Versuchen Sie, die Kalkulation wie gezeigt zu wiederholen und klicken Sie auf dem Tabellen Tab [P] um die erschaffene Wahrscheinlichkeitsdichtefunktions-Ergebnisse anzusehen. In diesem Beispiel, ist die binomiale Verteilung mit einem beginnenden Inputsatz von Probe = 20, Erfolgswahrscheinlichkeit = 0.5, und Zahl der erfolgreichen Proben X = 10, wo Erfolgswahrscheinlichkeit wechseln darf von 0., 0.25, ..., 0.50 und wird als Zeilenvariable aufgezeigt, und auch Zahl der erfolgreichen Proben von 0, 1, 2, ..., 8, wechseln darf, und wird als Säulenvariable aufgezeigt. Die PDF ist ausgewählt, und infolgedessen zeigen die Ergebnisse in der Tabelle die Wahrscheinlichkeit der vorgegebenen Abläufe. Zum Beispiel, die Wahrscheinlichkeit genau 2 Erfolge zu erzielen wenn 20 Proben laufen wobei jede Probe eine 25% Erfolgschance hat, ist 0.0669 oder 6.69% Wahrscheinlichkeit.

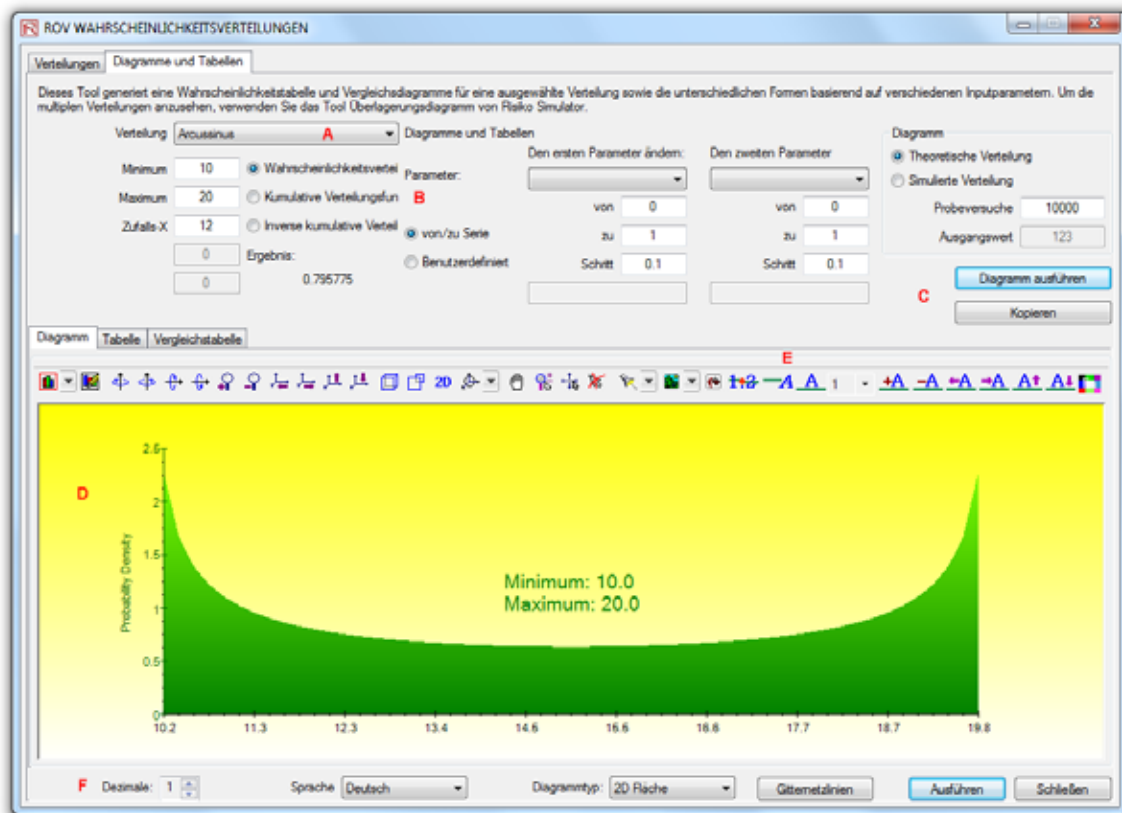


Bild 5.49—ROV Wahrscheinlichkeitsverteilung (PDF und CDF Charts)

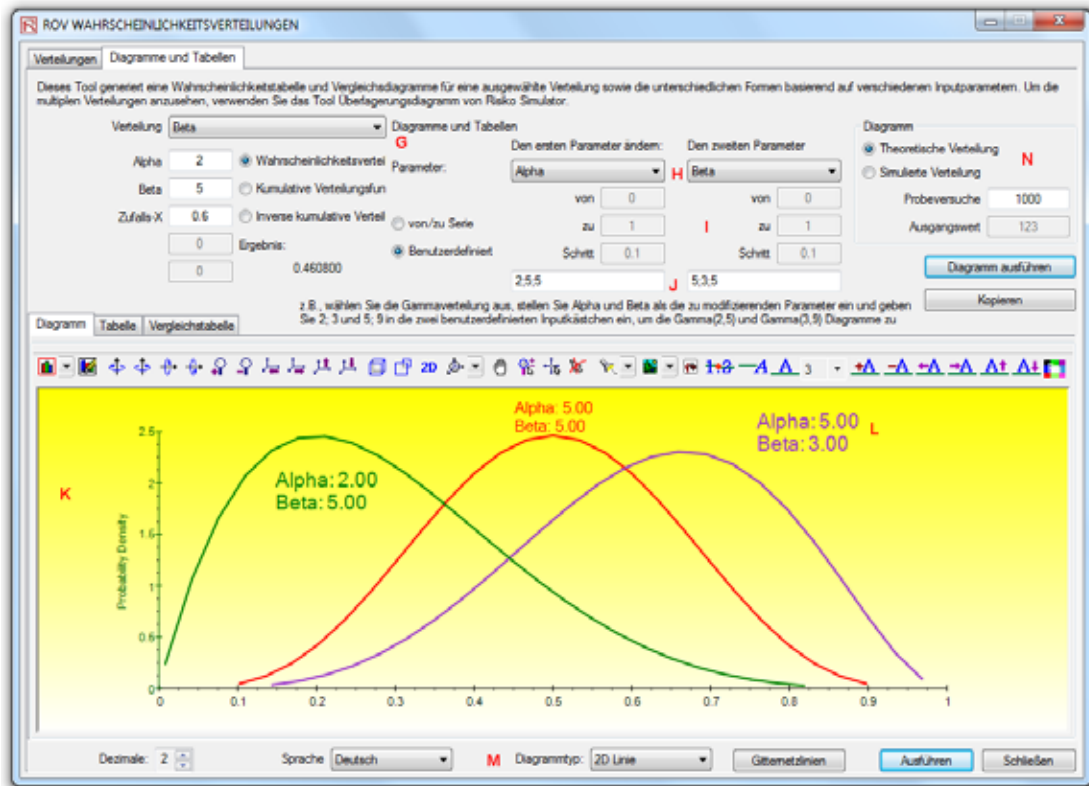


Bild 5.50—ROV Wahrscheinlichkeitsverteilung

ROV WAHRSCHEINLICHKEITSVERTEILUNGEN

Verteilungen Diagramme und Tabellen

Dieses Tool generiert eine Wahrscheinlichkeitstabelle und Vergleichsdiagramme für eine ausgewählte Verteilung sowie die unterschiedlichen Formen basierend auf verschiedenen Inputparametern. Um die multiplen Verteilungen anzusehen, verwenden Sie das Tool Überlagerungsdiagramm von Risiko Simulator.

Verteilung: **Binomiale** Diagramme und Tabellen

Probeversuche: 20 Wahrscheinlichkeitsverteilung

Wahrscheinlichkeit: 0.5 Kumulative Verteilungsfunktion

Zufalls-X: 10 Inverse kumulative Verteilung

Ergebnis: 0.176197

Den ersten Parameter ändern: Wahrscheinlichkeit von 0.2 zu 0.5 Schritt: 0.05

Den zweiten Parameter ändern: Zufalls-X von 0 zu 8 Schritt: 1

Diagramm: Theoretische Verteilung (ausgewählt) Simulierte Verteilung

Probeversuche: 1000 Ausgangswert: 123

Tabelle ausführen Kopieren

Diagramm Tabelle Vergleichstabelle

Reihenvariable: Wahrscheinlichkeit Spaltenvariable: Zufalls-X Typ: Wahrscheinlichkeitsverteilungsfur

N	0.0000	1.0000	2.0000	3.0000	4.0000	5.0000	6.0000	7.0000	8.0000
1	0.2000	0.0115	0.0576	0.1369	0.2054	0.2182	0.1746	0.1091	0.0545
2	0.7500	0.0032	0.0211	0.0669	0.1339	0.1897	0.2023	0.1686	0.1124
3	0.3000	0.0008	0.0068	0.0278	0.0716	0.1304	0.1789	0.1916	0.1643
4	0.3500	0.0002	0.0020	0.0100	0.0323	0.0738	0.1272	0.1712	0.1844
5	0.4000	0.0000	0.0005	0.0031	0.0123	0.0350	0.0746	0.1244	0.1659
6	0.4500	0.0000	0.0001	0.0008	0.0040	0.0139	0.0365	0.0746	0.1221
7	0.5000	0.0000	0.0000	0.0002	0.0011	0.0046	0.0148	0.0370	0.0739

Dezimal: 4 Sprache: Deutsch Ausführen Schließen

Bild 5.51—ROV Wahrscheinlichkeitsverteilung (Verteilungstabellen)

ROV BizStats

Dieses neue ROV BizStats Tool ist ein sehr starkes und schnelles Modul im Risiko Simulator, dass für den Ablauf von Betriebsstatistik und analytischen Modellen auf Ihren Daten angewendet wird, und erfasst mehr als 130 Betriebsstatistik und analytische Modelle (Bild 5.52-5.55). Folgendes liefert ein paar Schritte für das schnelle Loslegen des Modul-Ablaufes und Details über jedes Element in der Software.

Ablauf: Procedure:

- ROV BizStats auf *Risiko Simulator* | *ROV BizStats* laufen lassen, *Beispiel* anklicken um Stichprobendaten und Musterprofil zu laden [A] oder Ihre Daten eintippen oder kopieren/einfügen ins Datenraster [D] (Bild 5.52). Sie können Ihre eigene Notizen oder variable Namen in die ersten Notizen-Reihe hinzufügen [C].
- Relevantes Model um in Stufe 2 laufen zu lassen auswählen [F], und mit Anwendung der Musterdaten Input-Einstellungen [G], die relevanten Variablen eingeben [H]. Variablen separieren für den gleichen Parameter unter Verwendung von Semikola, und eine neue Zeile benutzen (Eingabe drücken um eine neue Zeile zu erstellen) für verschiedene Parameter.
- *Ausführen* [I] klicken um die Ergebnisse zu errechnen [J]. Sie können jegliche relevante analytische Ergebnisse, Charts oder Statistiken aus den verschiedenen Tabs in Stufe 3 ansehen.
- Wenn nötig, können Sie einen Modelname erstellen um in das Profil in Stufe 4 abzuspeichern [L]. Mehrere Modelle können in dem gleichen Profil abgespeichert werden. Bestehende Modelle können editiert oder gelöscht werden [M], neu angeordnet in der Reihenfolge des Auftretens [N], und alle Änderungen können in einem einzelnen Profil mit der Dateinamenerweiterung *bizstats abgespeichert werden [O].

Notizen:

- Die Daten-Rastergröße kann im Menü eingerichtet werden, wo das Raster bis zu 1.000 variable Säulen mit 1 Million Datenreihen pro Variable unterbringen kann. Das Menü erlaubt Ihnen die Spracheinstellungen und Dezimalstellungen für Ihre Daten zu ändern.
- Um loszulegen, ist es immer eine gute Idee die Musterdatei [A], die komplett mit einigen Daten und vorgefertigten Modellen kommt [S] zu laden. Sie können den Doppelklick anwenden um jede dieser Modelle ablaufen zu lassen, und die Ergebnisse werden im Ergebnisfeld gezeigt [J], mitunter als Chart oder Model-Statistiken [T/U]. Mit der Anwendung dieses Beispiel- Files können Sie nun sehen wie die Inputparameter [H], basierend auf der Modelbeschreibung, eingegeben sind [G], und Sie können jetzt fortfahren um Ihre eigene selbsterstellte Modelle zu kreieren.
- *Variable Header* [D] anklicken um eine oder multiple Variablen gleichzeitig auszuwählen,

dann rechtsklicken um die ausgewählte Variablen zuzufügen, löschen, kopieren, einfügen oder visualisieren [**P**].

- Modelle können auch unter Anwendung einer Befehlskonsole eingegeben werden [**V/W/X**]. Um zu sehen wie dies funktioniert, doppelklicken um ein Modell laufen zu lassen [**S**] und zur Befehlskonsole gehen [**V**]. Sie können das Modell replizieren oder Ihr eigenes herstellen und, wenn bereit, **Ausführen Befehl** klicken [**X**]. Jede Zeile in der Konsole repräsentiert ein Modell und seine relevante Parameter.
- Das gesamte *.bizstats Profil (wo Daten und multiple Modelle erstellt und gespeichert sind) kann direkt in XML editiert werden [**Z**] mit Öffnung der XML Editor aus dem File Menü. Profil-änderungen können hier programmatisch gemacht werden, und werden in Kraft treten sobald das File gespeichert ist.

Tipps

- Spaltenkopf(köpfe) des Datenrasters anklicken um die ganze Spalte(n) oder Variable(n) auszuwählen, dann, einmal ausgewählt, können Sie auf dem Kopf rechtsklicken um die Spalte automatisch anzupassen (Auto Fit), Daten Ausschneiden, Kopieren, Löschen oder Einfügen. Sie können auch multiple Spaltenköpfe anklicken und auswählen um multiple Variablen auszuwählen, dann rechtsklicken und Visualize Visualisieren auswählen um die Daten aufzuzeichnen.
- Wenn eine Zelle einen großen Wert hat der nicht vollständig dargestellt wird, klicken diese Zelle an, fahren Sie Ihre Maus über diese Zelle und Sie werden eine Pop-up Kommentar sehen, der den vollständigen Wert zeigt. Oder einfach die variable Spalte in der Größe anpassen (die Spalte ziehen um sie breiter zu machen, auf dem Spaltenrand doppelklicken um die Spalte automatisch anzupassen, oder auf dem Spaltenkopf rechtsklicken und Auto Fit auswählen.
- Die Oben- Unten- Links- und Rechts- Tasten anwenden um in das Raster heranzuziehen, oder die Home und End Tasten auf der Tastatur anwenden um ganz nach links und ganz nach rechts in eine Reihe zu bewegen. Sie können auch Kombinationstasten anwenden wie: Strng+Home um in die Zelle oben links zu springen, Strng+End um in die Zelle unten rechts, Shift+Oben/Unten um einen bestimmten Bereich auszuwählen, usw.
- Sie können Kurzmitteilungen für jede Variable auf der Mitteilungsreihe eingeben. Denken Sie daran, Ihre Mitteilungen kurz und schlicht zu gestalten.
- Probieren Sie die verschiedenen Chart Icons auf der Visualisieren Taste aus um das Erscheinungsbild der Charts zu verändern (z.B. rotieren, verschieben, zoomen, Farben ändern, Legende zufügen und so weiter).
- Die Kopieren Taste wird angewendet um die Ergebnisse, Charts, und Statistik Tabulatoren in Stufe 3 zu kopieren nachdem ein Modell ausgeführt wurde. Wenn keine Modelle ausgeführt wurden, wird die Kopieren Funktion bloß eine leere Seite kopieren.
- Die Bericht Taste wird nur dann laufen wenn sich in der Stufe 4 gespeicherte Modelle

befinden, oder wenn sich Daten im Raster befinden, sonst wird der generierte Bericht leer sein. Sie werden auch Microsoft Excel installiert haben müssen um die Datenextrahierung und Ergebnisberichte ausführen zu können, und Microsoft Power Point verfügbar um die Chart Berichte auszuführen.

- Im Zweifelsfall wie ein bestimmtes Modell oder eine statistische Methode ausgeführt wird, starten Sie das Beispiel Profil und überprüfen Sie wie die Daten in Stufe 1 eingerichtet sind, oder wie die Input Parameter in Stufe 2 eingegeben wurden. Sie können diese als Einstiegshilfen und Mustervorlagen für Ihre eigene Daten und Modelle anwenden.
- Die Sprache kann im Sprache Menü geändert werden. Beachten Sie, dass gegenwärtig 10 Sprachen in der Software verfügbar sind und weitere später hinzugefügt werden. Allerdings, werden manchmal bestimmte begrenzte Ergebnisse noch in Englisch gezeigt.
- Sie können das Erscheinungsbild der Modellliste in Stufe 2 ändern in dem Sie die Ansicht Dropliste verändern. Sie können sie Modelle alphabetisch, kategorisch und nach Daten Input Erfordernissen auflisten – beachten Sie, dass es in bestimmten Unicode Sprachen (z.B. Chinesisch, Japanisch und Koreanisch) keine alphabetische Anordnung gibt und deshalb die erste Option nicht verfügbar sein wird.
- Die Software kann verschiedene regionale, dezimale und numerische Einstellungen bearbeiten (z.B. einhunderttausend Dollar und fünfzig Cent kann als 1,000.50 oder 1.000,50 oder 1'000,50 und so weiter geschrieben werden). Die dezimale Einstellungen können in ROV BizStats' Menü **Daten | Dezimale Einstellungen** eingestellt werden. Allerdings, bitte im Zweifelsfall die regionale Einstellungen des Computers auf English USA umändern und die Voreinstellung North America 1,000.50 in ROV BizStats behalten (es ist garantiert, dass diese Einstellung mit ROV BizStats und den vorgegebenen Beispielen zusammenarbeitet).

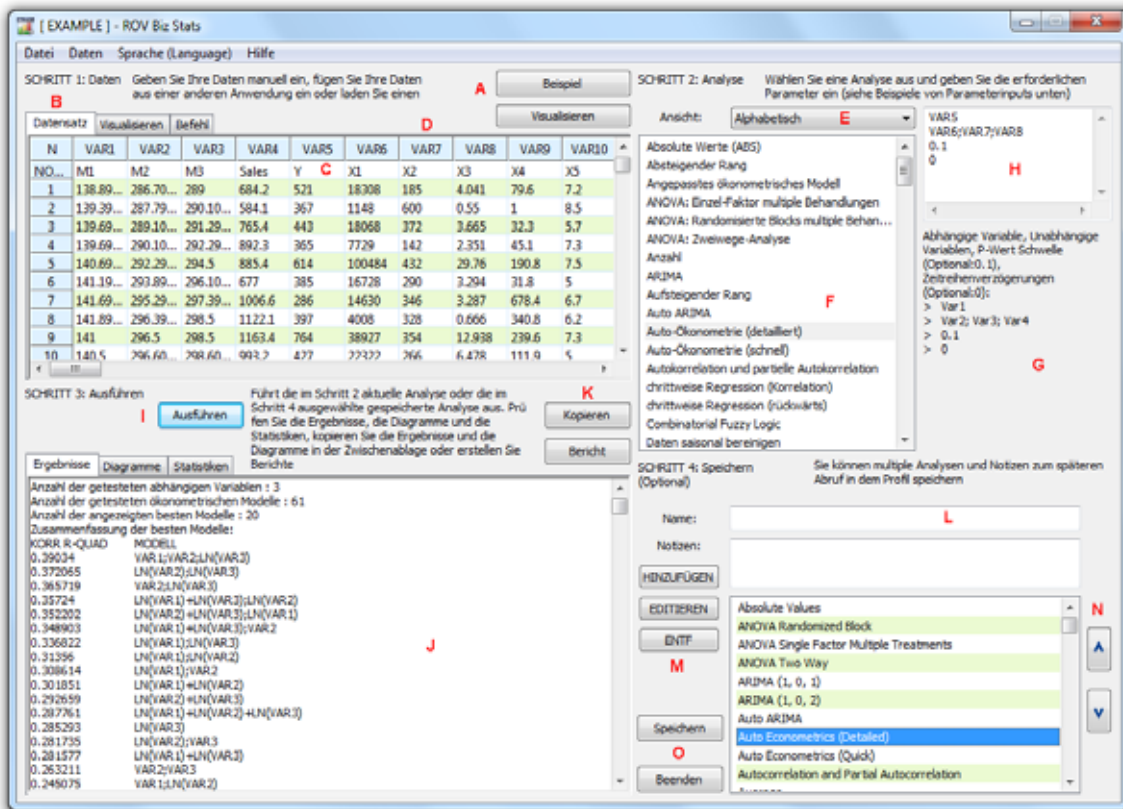


Bild 5.52—ROV BizStats (Statistische Analyse)

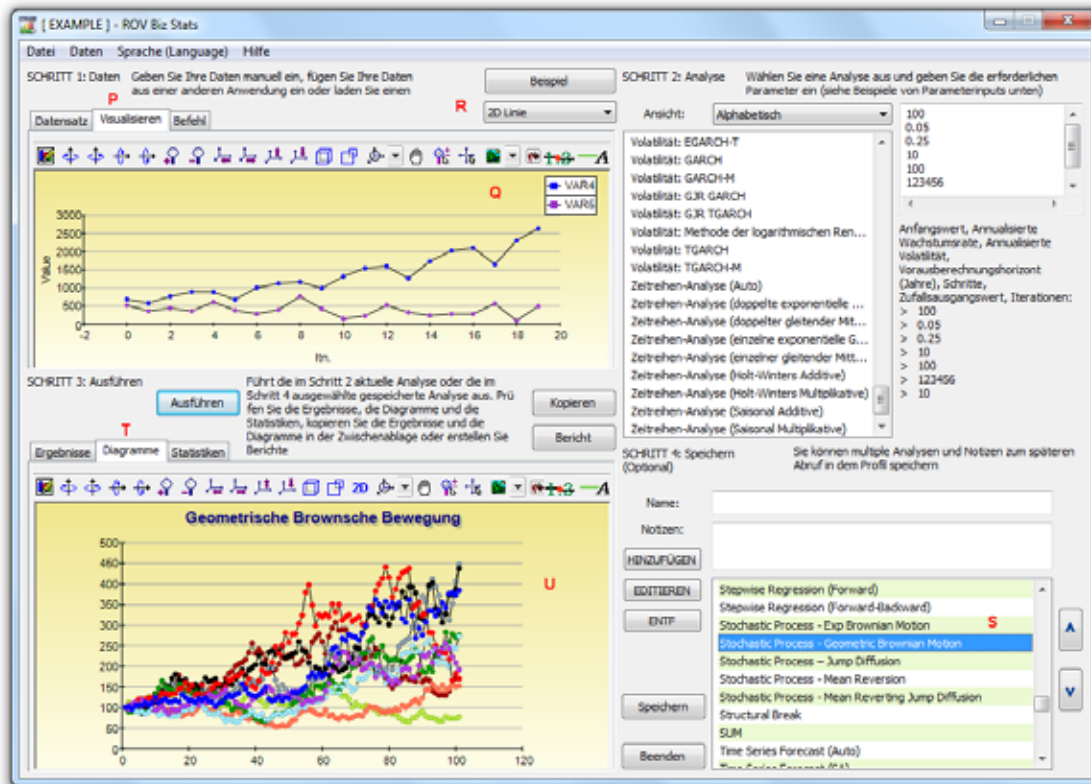


Bild 5.53—ROV BizStats (Datenvisualisierung und Ergebnis Charts)

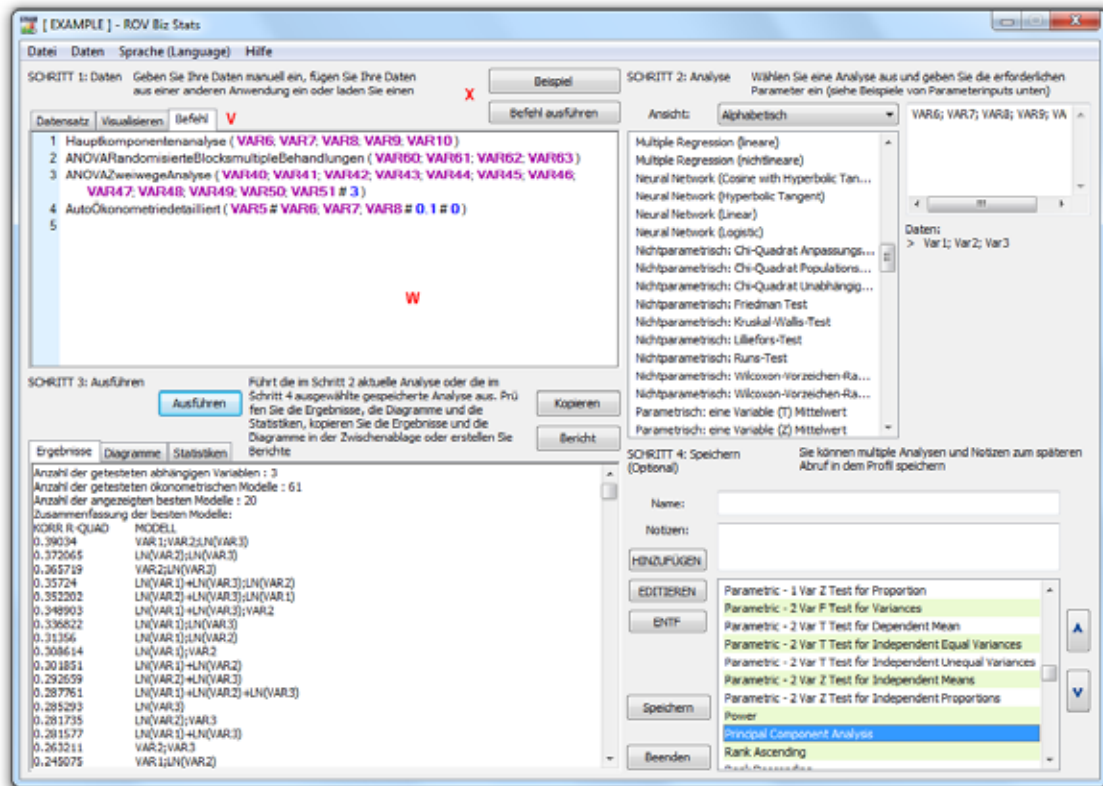


Bild 5.54—ROV BizStats (Befehlskonsole)

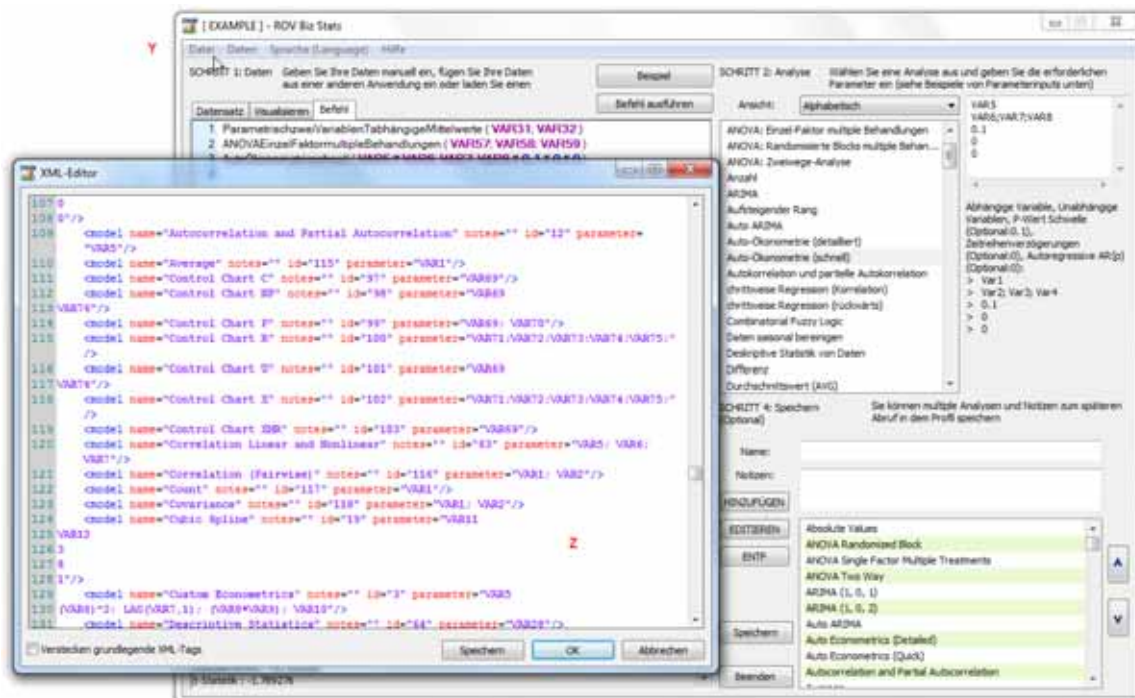


Bild 5.55—ROV BizStats (XML Editor)

Neuronales Netzwerk und Kombinatorische Fuzzy Logic Voraussagen

Methodologien

Der Begriff Neuronales Netzwerk wird oft benutzt um auf ein Netzwerk oder Kreis biologischer Neuronen zu verweisen, während moderne Verwendung der Terminologie oft auf künstliche neurale Netze verweist die aus künstlichen Neuronen oder Knotenpunkte bestehen, die in einer Software Umgebung wiederhergestellt werden. Solche Netzwerke versuchen die Neuronen im Menschlichen Gehirn Denkweisen und Mustererkennung nachzuahmen, und in unserer Situation, Muster erkennen für die Zwecke der voraussagenden Zeitreihendaten. Im **Risiko Simulator**, befindet sich die Methodik innerhalb des **ROV BizStats** Moduls im **Risiko Simulator | ROV BizStats | Neuronales Netzwerk** sowie im **Risiko Simulator | Voraussagen | Neuronales Netzwerk**. Bild 5.56 zeigt die Neuronale Voraussagen Methodik.

Ablauf Procedure

- **Risiko Simulator | Voraussagen | Neuronales Netzwerk** anklicken.
- Beginnen Sie mit der manuellen Einspeisung der Daten oder die Einfügung einige Daten aus der Zwischenablage (z.B. einige Daten von Excel auswählen und kopieren, dieses Tool starten, und die Daten mit klicken auf der **Einfügen** Taste die Daten einfügen).
- Auswählen, ob Sie ein **Linear or Nicht linear** neuronales Netzwerk-Modell laufen lassen möchten, die erwünschte Zahl der **Voraussage Perioden** (z.B., 5), die Zahl der verdeckten **Schichten** im neuronalen Netzwerk (z.B., 3), und Zahl der **Testperioden** (z.B., 5) eingeben.
- **Ausführen**Taste anklicken um die Analyse durchzuführen und die errechnete Ergebnisse und Charts zu überprüfen. Sie können auch die Ergebnisse und das Chart in die Zwischenablage **kopieren** und in eine andere Anwendersoftware einfügen.

Beachten Sie, dass die Zahl der verdeckten Schichten im Netzwerk ein Input Parameter ist, die nötigerweise mit Ihren Daten kalibriert werden soll. Typischerweise, um so komplizierte der Datenmuster, um so höher die benötigten Zahl der verdeckten Schichten und um so länger die Dauer der Berechnung. Es wird empfohlen, dass Sie mit 3 Schichten anfangen Die Testperiode ist lediglich die Zahl der Datenpunkte die in der abschließenden Kalibrierung des neuronalen Netzwerkmodells angewendet wurde, und wir empfehlen mindestens die gleiche Periodenzahl die Sie voraussagen möchten anzuwenden als Testperiode.

Hingegen ist der Begriff *Fuzzy Logik* von der Fuzzy Set Theorie abgeleitet, die mit approximativer statt akkurater Argumentation umgeht—im Gegensatz zu *Crisp Logik* wo Binärsätze binäre Logik haben, können Fuzzy Logik Variablen einen Wahrheitswert haben, der sich zwischen 0 und 1 bewegt und ist nicht auf die zwei Wahrheitswerte der klassischen Aussagenlogik beschränkt. Dieses fuzzy Wägungsschema wird zusammen mit einer kombinatorischen Methode angewendet um Zeitreihenprognosen Ergebnisse im **Risiko Simulator**

hervorzubringen wie in Bild 5.57 erklärt, und ist am meisten geeignet wenn es auf Zeitreihendaten mit Saisonalität und Trend angewendet wird. Diese Methodologie befindet sich innerhalb des *ROV BizStats* Modul im *Risiko Simulator* bei [Risiko Simulator | ROV BizStats | Kombinatorische Fuzzy Logik](#) sowie in [Risiko Simulator | Voraussagen | Kombinatorische Fuzzy Logik](#).

Ablauf

- [Risiko Simulator | Voraussagen | Kombinatorische Fuzzy Logik](#) anklicken.
- Beginnen Sie mit der manuellen Einspeisung der Daten oder die Einfügung einige Daten aus der Zwischenablage (z.B. einige Daten von Excel auswählen und kopieren, dieses Tool starten, und die Daten mit klicken auf der *Einfügen* Taste die Daten einfügen).
- Die Variable für der Sie die Analyse durchführen möchten aus der Drop-Down-Liste auswählen, und in der Saisonalitätsperiode (z.B., 4 für vierteljährliche Daten, 12 für monatliche Daten, usw.) und die erwünschte Zahl der *Voraussagen Perioden* (z.B., 5) eingeben.
- *Ausführen* anklicken um die Analyse durchzuführen und die errechnete Ergebnisse und Charts zu überprüfen. Sie können auch die Ergebnisse und Charts in die Zwischenablage *Kopieren* und in eine andere Software Applikation einfügen.

Beachten Sie, dass weder neuronale Netzwerke noch Fuzzy Logik Techniken als gültige und zuverlässige Methoden im Bereich der Umsatzvoraussage auf einer strategische, taktische oder betriebliche Ebene etabliert sind. Viel Forschung in diesen erweiterten Bereichen der Voraussagen ist noch erforderlich. Trotzdem liefert *Risiko Simulator* die Grundlagen dieser zwei Techniken für die Zwecke der Ausführungen von Zeitreihenprognosen. Wir empfehlen Ihnen jede diese Techniken nicht isoliert anzuwenden, sondern mit anderen *Risiko Simulator* Voraussagen Methodologien zu kombinieren, um robustere Modelle zu erstellen.

Vorausberechnung nach neuronales Netzwerk

SCHRITT 1: Daten Geben Sie Ihre Daten manuell ein, fügen Sie Ihre Daten aus einer anderen Anwendung ein oder laden Sie einen Beispieldatensatz mit Analyse Einfügen

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								

SCHRITT 2: Den auszuführenden Analysentyp, die auszuführende Variable und die auszuführende Vorausberechnungsperiode auswählen

☐ Kosinus mit hyperbolischer Tangente
☒ Hyperbolische Tangente
☐ Linear(e)
☐ Logistisch(e)

Ebenen:
 Satztest wird durchgeführt:
 Vorausberechnungsperioden:

☒ Mehrphasige Optimierung anwenden

Ergebnisse ☒ Diagramme

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

Bild 5.56—Neuronale Netzwerk Voraussage

Vorausberechnung nach kombinatorischer Fuzzylogik

SCHRITT 1: Daten Geben Sie Ihre Daten manuell ein, fügen Sie Ihre Daten aus einer anderen Anwendung ein oder laden Sie einen Beispieldatensatz mit Analyse Einfügen

N	VAR1	VAR2	VAR3	VAR4	VAR5	VAR6	VAR7	VAR8	VAR9	VAR10
NOT...	FUZZY									
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									

SCHRITT 2: Die erforderlichen Inputs eingeben und die vorauszuberechnende Variable auswählen Deutsch

Saisonalität: VAR1 Kopieren

Vorausberechnungsperioden: 4 Ausführen

Ergebnisse Diagramme

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

RESULTS
Forecast Fit
* indicates negative values

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

Bild 5.57—Fuzzy Logik Zeitreienvoraussage

Optimierer Goal Seek

Das *Goal Seek* Tool ist ein Suchalgorithmus das um die Lösung einer Einzelvariablen innerhalb eines Modells zu finden angewendet wird. Wenn Sie das Ergebnis wissen das Sie von einer Formel oder einem Modell haben wollen, sind aber nicht sicher welchen Inputwert die Formel benötigt um das Ergebnis zu erzielen, nutzen Sie die Goal Seek Funktion Beachten Sie, dass *Goal Seek* mit nur einem Variablen Input Wert funktioniert. Wenn Sie mehr als einen Input Wert annehmen möchten, nutzen Sie *Risiko Simulator* erweiterte Optimierungs-Routinen. Bild 5.58 zeigt wie *Goal Seek* für ein einfaches Modell angewandt wird.

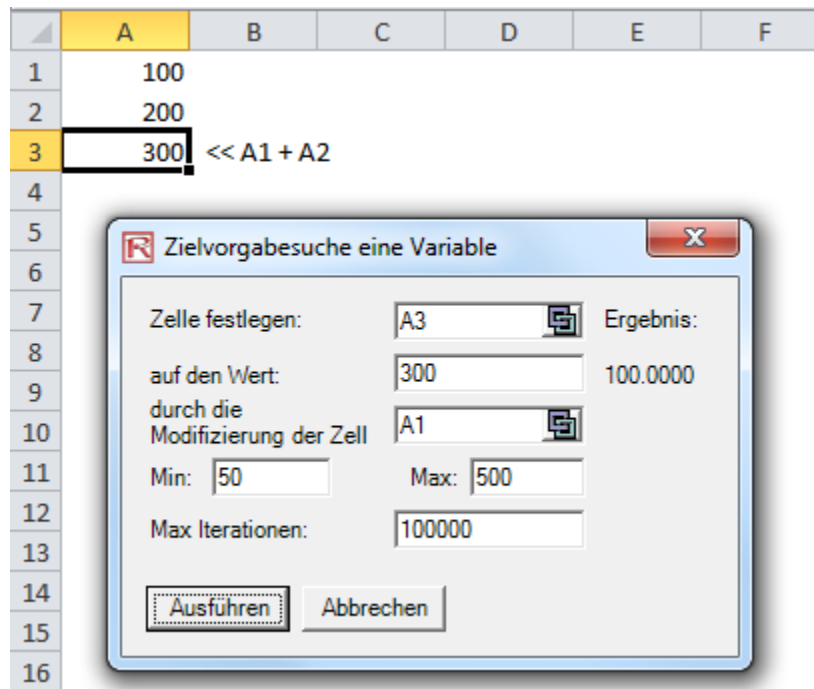


Bild 5.58—Goal Seek

Einzel Variable Optimierer

Das *Einzel Variable Optimierer* Tool ist ein Suchalgorithmus das um die Lösung einer Einzelvariablen innerhalb eines Modells zu finden angewendet wird, genau so wie die Goal Seek Routine die vorher erörtert wurde Wenn Sie das maximal oder minimal möglichste Ergebnis eines Modells wollen, sind aber nicht sicher welchen Input-Wert die Formel benötigt um dieses Ergebnis zu erzielen, nutzen Sie die *Risiko Simulator* | *Tools* | *Einzel Variable Optimierer* Funktion (Bild 5.59). Beachten Sie, dass dieses Tool sehr schnell läuft, ist aber nur anwendbar um einen Variable Input zu finden. Wenn Sie mehr als einen Input Wert einsehen möchten, *Risiko Simulator* erweiterte Optimierungs-Routinen anwenden. Beachten Sie, dass dieses Tool im *Risiko Simulator* einbezogen ist, weil wenn Sie eine schnelle Optimierungsberechnung für eine Einzel-Entscheidungsvariable brauchen, liefert dieses Tool genau diese Leistungsfähigkeit ohne ein

Optimierungsmodell mit Profilen, Simulationsannahmen, Entscheidungsvariablen, Objektiven und Randbedingungen einrichten zu müssen.

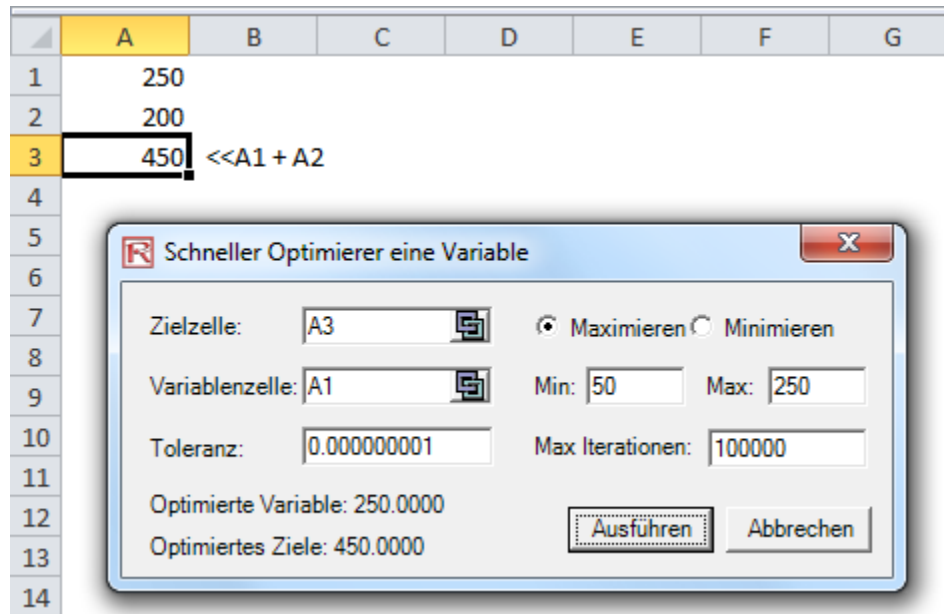


Bild 5.59—Einzel Variable Optimierer

Genetische Algorithmus Optimierung

Genetische Algorithmen gehören zu der größeren Kategorie der evolutionären Algorithmen, die Lösungen von Optimierungsproblemen generieren mit der Anwendung von Techniken die von der natürlichen Evolution wie Vererbung, Mutation, Selektion und Kreuzung inspiriert sind. Genetischer Algorithmus ist eine Suchheuristik die der Prozess der natürlichen Evolution nachahmt und routinemäßig angewendet wird um nützliche Lösungen zur Optimierung und Suchprobleme zu generieren.

Der genetische Algorithmus ist im [Risiko Simulator | Tools | Genetischer Algorithmus](#) abrufbar(Bild 5.60). Die Kalibrierung des Modell-Inputs sollte sorgfältig durchgeführt werden, weil die Ergebnisse ziemlich empfindlich auf die Inputs sein werden (vorgegebene Inputs stehen als allgemeine Anleitung für die gängigste Input-Ebenen zur Verfügung), und es wird empfohlen, die [Gradient Suchtest](#) Option auszuwählen. eine robustere Ergebnisreihe zu erzielen. (Um anzufangen, können Sie diese Option deaktivieren ([die Markierung dieser Option aufheben](#)), dann diese Auswahl selektieren, die Analyse nochmals laufen lassen und die Ergebnisse vergleichen

Notizen Notes

In vielen Problemen genetische Algorithmen könnten eine Tendenz zeigen, sich lokale Optima oder sogar arbiträre Punkte eher als das globale Optimum des Problems anzunähern. Dies bedeutet, dass es nicht weiß wie kurzfristige Eignung zu opfern ist um längerfristige Eignung zu erreichen. Für spezifische Optimierungsprobleme und Problemfälle, könnten vielleicht andere Optimierungs-Algorithmen bessere Lösungen als genetische Algorithmen entdecken (in Anbetracht der gleichen Menge Berechnungszeit). Es ist deshalb zu empfehlen, dass Sie als erstes den *Genetischer Algorithmus* durchführen und dann, um die Robustheit des Modelles zu prüfen, mit Vorgabe der *Gradient Suchtest Anwenden* Option (Bild 5.60) wiederholen Dieser Gradient Suchtest wird versuchen, Kombinationen der traditionellen Optimierungstechniken mit Genetischen Algorithmus Methoden durchzuführen und die bestmögliche Lösung wiederzubringen Zum Schluss, sofern nicht eine spezifische theoretische Notwendigkeit *Genetischer Algorithmus* zu nutzen besteht, empfehlen wir die Nutzung des *Risiko Simulators Optimierung* Modul, welches Ihnen erlaubt mehr fortgeschrittene risikobasierte dynamisch und stochastische Optimierungsroutinenen für robustere Ergebnisse auszuführen.

Genetischer Algorithmus

Zielzelle: ☒ Maximieren ☐ Minimieren

Variablen:

Zelle	Min	Max
-------	-----	-----

Bedingungen:

Zelle	Min	Max
-------	-----	-----

Max Iterationen: Mutationsrate:

Bevölkerungsgröße: Diversität:

Überkreuzungsrate: Elitismus:

Überkreuzung: Unverändert:

☐ Gradientensuchetest anwenden

Ergebnis:

Bild 5.60—Genetischer Algorithmus

Modul ROV Entscheidungsbaum

ROV Entscheidungsbaum (Bild 5.61) wird verwendet, um Entscheidungsbaummodelle zu erstellen und bewerten. Zusätzliche fortgeschrittene Methodologien und Analytiken sind ebenfalls enthalten:

- Entscheidungsbaummodelle
- Monte-Carlo-Risikosimulation
- Sensibilitätsanalyse
- Szenarienanalyse
- Bayessche-Analyse (gesamte und a posteriori Wahrscheinlichkeitsaktualisierung)
- Erwartungswert der Information
- MINIMAX
- MAXIMIN
- Risikoprofile

Es folgen einige Kurzeinstiegstipps und -prozeduren für die Verwendung dieses intuitiven Tools:

- In diesem Modul stehen 11 lokale Sprachen zur Verfügung und man kann die aktuelle Sprache im Sprach-Menü ändern.
- Die Funktionen *Optionsknoten eingeben* oder *Terminalknoten eingeben* erfolgen, indem Sie zuerst alle vorhandenen Knoten auswählen und dann auf das Optionsknotensymbol (Viereck) oder das Terminalknotensymbol (Dreieck) klicken; Sie können aber auch die Funktionen im Menü *Eingabe* verwenden.
- Ändern Sie die Eigenschaften von einzelnen *Optionsknoten* oder *Terminalknoten* durch das Doppelklicken auf einem Knoten. Gelegentlich, wenn Sie auf einem Knoten klicken, werden auch alle nachfolgenden Kindknoten ausgewählt (dieses erlaubt es Ihnen, den gesamten Baum ab diesem ausgewählten Knoten zu verschieben). Wenn Sie nur diesen Knoten auswählen möchten, könnte es sein, dass Sie auf den leeren Hintergrund und dann zurück auf diesem Knoten klicken müssen, um diesen Knoten einzeln auszuwählen. Im Weiterem können Sie, abhängig von den aktuellen Einstellungen, einzelne Knoten oder den gesamten Baum ab dem ausgewählten Knoten verschieben (mit einem Doppelklick oder gehen Sie im Menü *Editieren* und wählen Sie *Knoten einzeln verschieben* oder *Knoten zusammen verschieben*).
- Im Folgenden sind einige Kurzbeschreibungen der Elemente, die man in der Bedienoberfläche der Knoteneigenschaften anpassen und konfigurieren kann. Am einfachsten ist es, verschiedene Einstellungen für jedes der folgenden Elementen auszuprobieren, um ihre Auswirkungen auf dem Strategiebaum zu sehen:
 - o *Name*. Name, der über dem Knoten angezeigt wird.
 - o *Wert*. Wert, der über dem Knoten angezeigt wird.
 - o *Excel-Link*. Verknüpft den Wert aus der Zelle einer Excel-Tabellenkalkulation.
 - o *Notizen*. Notizen können über oder unter einen Knoten eingefügt werden.
 - o *Im Modell anzeigen*. Zeigt beliebigen Kombinationen von Namen, Werten und Notizen an.
 - o *Lokale Farbe gegen Globale Farbe*. Die Knotenfarben können lokal für einen Knoten oder global geändert werden.
 - o *Beschriftung in der Form*. Text kann innerhalb der Form eingefügt werden (es ist möglich, dass Sie den Knoten verbreitern müssen, um einen längeren Text unterzubringen).

- o *Name des Verzweigungsereignisses.* Text kann auf der Verzweigung, die zu dem Knoten führt platziert werden, um auf das Ereignis, welches zu diesem Knoten führt zu weisen.
 - o *Reelle Optionen auswählen.* Einen spezifischen Typ von reeller Option kann dem aktuellen Knoten zugeordnet werden. Die Zuordnung reeller Optionen zu den Knoten erlaubt es dem Tool, eine Liste der erforderlichen Inputvariablen zu generieren.
- *Globale Elemente* sind alle anpassbar, einschließlich die Elemente *Hintergrund*, *Verbindungslinien*, *Optionsknoten*, *Terminalknoten* und *Textfelder* des Strategiebaums. Zum Beispiel, man kann die folgenden Einstellungen für jedes der Elemente ändern:
 - o *Schriftart-Einstellungen* für Namen, Wert, Notizen, Beschriftung, Ereignisnamen.
 - o *Knotengröße* (minimale und maximale Höhe und Breite).
 - o *Rahmen* (Linienstile, Breite und Farbe).
 - o *Schattierung* (Farben und Anwendung oder Nichtanwendung von einer Schattierung).
 - o *Globale Farbe.*
 - o *Globale Form.*
- Das Befehl *Datenanforderungsfenster anzeigen* des Menüs *Editieren* öffnet ein gedocktes Fenster auf der rechten Seite des Strategiebaums, sodass, wenn ein Optionsknoten oder Terminalknoten ausgewählt wird, die Eigenschaften dieses Knotens angezeigt und direkt aktualisiert werden können. Diese Funktion liefert eine Alternative zur Doppelklicken auf einem Knoten jedes Mal.
- *Beispielsdateien* stehen im Menü *Datei* zur Verfügung, um Sie bei der anfänglichen Erstellung von Strategiebäumen zu helfen.
- *Datei schützen* aus dem Menü *Datei* erlaubt die Verschlüsselung des Strategiebaums mit einer Passwortverschlüsselung bis zu 256-Bit. Geben Sie Acht, wenn eine Datei verschlüsselt wird; beim Verlust des Passworts, kann die Datei nicht mehr geöffnet werden.
- Die Funktion *Bildschirm erfassen* oder das Drucken des existierenden Modells wird im Menü *Datei* durchgeführt. Der erfasste Bildschirm kann dann in anderen Softwareanwendungen eingefügt werden.
- Die Funktionen *Hinzufügen*, *Duplizieren*, *Umbenennen* und *Ein Strategiebaum löschen* können durch das Doppelklicken der Registerkarte Strategiebaum oder im Menü *Editieren* ausgeführt werden.
- Sie können auch folgende Funktionen ausführen: *Dateiverknüpfung einfügen* und *Kommentar einfügen* auf jedem beliebigen Options- oder Terminalknoten, oder *Text einfügen* oder *Abbildung einfügen* überall im Hintergrund- oder Zeichenbereich.
- Von Ihrem Strategiebaum, können Sie *Existierende Stile ändern* oder *angepasste Stile verwalten und kreieren* (dies schließt die Spezifikationen für Größe, Form, Farbschemas und Schriftartgröße/Farbe des gesamten Strategiebaums).
- Die Funktionen *Entscheidungsknoten eingeben*, *Ungewissheitsknoten eingeben* oder *Terminalknoten eingeben* erfolgen, indem Sie zuerst alle vorhandenen Knoten auswählen und dann auf das Entscheidungsknotensymbol (Viereck), Ungewissheitsknotensymbol (Kreis) oder das Terminalknotensymbol (Dreieck) klicken; Sie können aber auch die Funktionen im Menü *Eingabe* verwenden.
- Ändern Sie die Eigenschaften von einzelnen Entscheidungsknoten, Ungewissheitsknoten oder Terminalknoten durch das Doppelklicken auf einem Knoten. Im Folgenden sind einige einmalige zusätzliche Elemente im Modul Entscheidungsbaum, die man in der Bedienoberfläche der Knoteneigenschaften anpassen und konfigurieren kann.

- o Entscheidungsknoten: *Angepasste Überschreibung* oder *Auto-Berechnung* des Werts auf einem Knoten. Die Option Auto-Berechnung ist als Standard eingestellt; wenn Sie auf *AUSFÜHREN* von einem abgeschlossenen Entscheidungsbaummodell klicken, werden die Entscheidungsknoten mit den Ergebnissen aktualisiert.
- o Ungewissheitsknoten: *Ereignisnamen*, *Wahrscheinlichkeiten* und *Simulationshypothesen* festlegen. Sie können Wahrscheinlichkeitsereignisnamen, Wahrscheinlichkeiten und Simulationshypothesen, nur nach dem die Ungewissheitsverzweigungen erstellt wurden.
- o Terminalknoten: *Manuelle Eingabe*, *Excel-Link* und *Simulationshypothesen* festlegen. Die Terminalereignispayoffs können manuell eingegeben oder mit einer Excel-Zelle verknüpft werden (z.B. wenn Sie ein großes Excel-Modell haben, welches das Payoff berechnet, können Sie das Modell zu der Output-Zelle dieses großen Modells verknüpfen) oder die Wahrscheinlichkeitsverteilungshypothesen zur Ausführung von Simulationen festlegen.
- *Knoteneigenschaftenfenster anzeigen* ist im Menü *Editieren* verfügbar; die Eigenschaften des ausgewählten Knotens werden sich aktualisieren, wenn ein Knoten ausgewählt wird.
- Das Modul Entscheidungsbaum enthält auch die folgenden fortgeschrittenen Analytiken:
 - o Monte-Carlo-Simulation Modellierung auf Entscheidungsbäume
 - o Bayessche Analyse zur Beschaffung von posterior Wahrscheinlichkeiten
 - o Erwartungswert der perfekten Information, MINIMAX und MAXIMIN Analyse, Risikoprofile und Wert der unvollständigen Information
 - o Sensibilitätsanalyse
 - o Szenarienanalyse
 - o Nutzenfunktionsanalyse

Simulationsmodellierung

Dieses Tool führt eine Monte Carlo-Risikosimulation auf den Entscheidungsbaum aus (Bild 5.62). Dieses Tool erlaubt es Ihnen, die Wahrscheinlichkeitsverteilungen als Inputhypothesen zur Ausführung von Simulationen einzustellen. Sie können entweder eine Hypothese für den ausgewählten Knoten festlegen oder auch eine neue Hypothese festlegen und diese neue Hypothese (oder die zuvor erstellte Hypothesen) in einer numerischen Gleichung oder Formel verwenden. Zum Beispiel können Sie eine neue Hypothese genannt 'Normal' festlegen (z.B. eine Normalverteilung mit einem Mittelwert von 100 und einer Standardabweichung von 10) und eine Simulation in dem Entscheidungsbaum ausführen oder diese Hypothese in einer Gleichung wie $(100 * \text{Normal} + 15.25)$ verwenden. Erstellen Sie Ihr eigenes Modell im Dialogfeld des numerischen Ausdrucks. Sie können einfache Berechnungen verwenden oder existierende Variablen Ihrer Gleichung hinzufügen, indem Sie auf die Liste der existierenden Variablen klicken. Auch können Sie nach Bedarf neue Variablen der Liste hinzufügen, diese dann entweder als numerische Ausdrücke oder als Hypothesen.

Bayessche Analyse

Dieses bayessches Analysetool (Bild 5.63) kann auf zwei beliebigen Ungewissheitsereignissen, die entlang einem Pfad verknüpft sind, ausgeführt werden. Zum Beispiel: im Beispiel rechts sind

die Ungewissheiten A und B verknüpft, wobei Ereignis A als Erstes in der Zeitlinie und Ereignis B als Zweites eintritt. Das erste Ereignis A ist Marktforschung mit zwei Ausgängen (erfolgreich oder nicht erfolgreich). Das zweite Ereignis B ist Marktbedingungen, ebenfalls mit zwei Ausgängen (stark und schwach). Dieses Tool wird verwendet, um die gemeinsamen, die marginalen und die bayesschen a-posteriori aktualisierten Wahrscheinlichkeiten durch die Eingabe von vorherigen Wahrscheinlichkeiten und zuverlässigkeitsbedingten Wahrscheinlichkeiten zu berechnen. Man kann aber auch Zuverlässigkeitswahrscheinlichkeiten berechnen, wenn man a-posteriori aktualisierte bedingte Wahrscheinlichkeiten besitzt. Wählen Sie die relevante gewünschte Analyse weiter unten aus und klicken Sie auf 'Beispiel laden'. So können Sie Folgendes visualisieren: die Musterinputs, die der ausgewählten Analyse entsprechen, die im Raster rechts angezeigten Ergebnisse und welche Ergebnisse als Inputs in dem Entscheidungsbaum in der Abbildung verwendet sind.

Quick Procedures

- SCHRITT 1: Den Namen des ersten und den Namen des zweiten Ungewissheitsereignisses eingeben und die Anzahl der Wahrscheinlichkeitsereignisse (Naturzustände oder Ausgänge) für jedes Ereignis auswählen.
- SCHRITT 2: Die Namen aller Wahrscheinlichkeitsereignisse oder Ausgänge eingeben.
- SCHRITT 3: Die Vorwahrscheinlichkeiten und die bedingten Wahrscheinlichkeiten für jedes Ereignis oder Ausgang des zweiten Ereignisses eingeben. Die Wahrscheinlichkeiten müssen sich auf 100% summieren.

Erwartungswert der perfekten Information, Minimax und Maximin Analyse, Risikoprofile und Wert der unvollständigen Information

Dieses Tool berechnet den Erwartungswert der perfekten Information (EWPI), die Minimax- und Maximin-Analyse und auch das Risikoprofil und den Wert der unvollständigen Information (Bild 5.64). Dabei geben Sie die Anzahl der Entscheidungsverzweigungen oder -strategien unter Berücksichtigung (z.B. der Bau einer großen, mittleren oder kleinen Anlage) und die Anzahl der ungewissen Ereignisse oder Zustände der Naturausgänge (z.B. guter Markt, schlechter Markt) und auch die erwarteten Payoffs unter jedes Szenario ein.

Der Erwartungswert der perfekten Information (EWPI). Das heißt, angenommen Sie hätten eine perfekte Voraussicht und wüssten genau was zu tun (wegen ein besseres Erkenntnis der probabilistischen Ausgänge durch Marktforschung oder anderen Mitteln), EWPI berechnet, ob es einen Mehrwert in so einer Art von Information gibt (sprich, ob Marktforschung einen Mehrwert hinzufügt) verglichen mit naiveren Schätzungen der probabilistischen Naturzustände. Um vorzugehen, geben Sie die Anzahl der Entscheidungsverzweigungen oder -strategien unter Berücksichtigung (z.B., der Bau einer großen, mittleren, kleinen Anlage) und die Anzahl von den ungewissen Ereignissen oder Zuständen der Naturausgänge (z.B., guter Markt, schlechter Markt) und auch die erwarteten Payoffs unter jedes Szenario ein.

Minimax (Minimierung des maximalen Regrets) und Maximin (Maximierung des minimalen Payoffs) sind zwei alternative Ansätze zur Ermittlung des optimalen Entscheidungspfads. Diese beiden Ansätze werden nicht oft verwendet, liefern jedoch zusätzliche Erkenntnisse über dem Entscheidungsfindungsprozess. Geben Sie die Anzahl der existierenden Entscheidungsverzweigungen oder -pfade (z.B. der Bau einer großen, mittleren oder kleinen Anlage) und die Ungewissheitsereignisse oder Naturzustände unter jeden Pfad (z.B. gute Wirtschaft gegen schlechte Wirtschaft) ein. Dann vervollständigen Sie die Payoff-Tabelle für die verschiedenen Szenarien und berechnen Sie die Minimax- und Maximin-Ergebnisse. Sie können auch auf Beispiel laden klicken, um eine Beispielsberechnung zu sehen.

Sensibilität

Die Sensibilitätsanalyse (Bild 5.65) der Inputwahrscheinlichkeiten wird durchgeführt, um ihre Auswirkung auf die Werte der Entscheidungspfade zu bestimmen. Wählen Sie unten zuerst einen zu analysierenden Entscheidungsknoten und wählen Sie dann aus der Liste ein zu testendes Wahrscheinlichkeitsereignis aus. Wenn es mehrere Ungewissheitsereignisse mit identischen Wahrscheinlichkeiten gibt, können diese entweder unabhängig voneinander oder gleichzeitig analysiert werden.

Die Sensibilitätsdiagramme zeigen die Werte der Entscheidungspfade unter variierenden Wahrscheinlichkeitsniveaus an. Die numerischen Werte werden in der Ereignistabelle angezeigt. Die Position der Crossoverlinien, wenn vorhanden, repräsentiert, bei welchen probabilistischen Ereignissen, ein bestimmter Entscheidungspfad über einen anderen dominiert.

Szenarientabellen

Szenarientabellen (Bild 5.66) können erstellt werden, um die Outputwerte bei einigen Änderungen zum Input festzustellen. Sie können einen oder mehrere Entscheidungspfade zum Analysieren (die Ergebnisse von jedem ausgewählten Pfad werden als eine separate Tabelle und ein separates Diagramm repräsentiert) und einen oder zwei Ungewissheits- oder Terminalknoten als Inputvariablen zur Szenarientabelle auswählen.

Quick Procedures

- Wählen Sie aus der Liste unten einen oder mehrere zu analysierende Entscheidungspfade aus.
- Wählen Sie ein oder zwei Ungewissheitsereignisse oder einen oder zwei Terminal-Payoffs zur Modellierung aus.
- Entscheiden Sie, ob Sie die Wahrscheinlichkeit oder den Payoff alleine oder alle identischen Wahrscheinlichkeiten/Payoffs gleichzeitig ändern möchten.
- Inputszenario-Variationsbreite eingeben.

Generation der Nutzenfunktion

Nutzenfunktionen (Bild 5.67), oder $U(x)$, werden gelegentlich anstelle der Erwartungswerte von Terminal-Payoffs in einem Entscheidungsbaum verwendet. Man kann $U(x)$ zweierlei entwickeln: Mit langwierigem und detailliertem Experimentieren eines jeden möglichen Ausganges oder mit einer exponentiellen Extrapolationsmethode (verwendet hier). Man kann sie für folgende Arten von Entscheidungsträgern je nach Bedarf modellieren: risikoscheu (Nachteile sind katastrophaler oder schmerzhafter als eine gleiche vorteilhafte Möglichkeit), risikoneutral (Vorteile und Nachteile sind gleichwertig interessant) oder risikofreudig (vorteilhafte Möglichkeit ist interessanter). Sie können dann den minimalen und maximalen Erwartungswert ihrer Terminal-Payoffs und die Anzahl der dazwischenstehenden Datenpunkte eingeben, um die Nutzenkurve und -tabelle zu erstellen.

Wenn Sie ein 50:50 gewagtes Unternehmen hätten, wobei Sie einerseits entweder $\$X$ verdienen oder $\$X/2$ verlieren würden oder andererseits nichts riskieren und einen $\$0$ Payoff erhalten würden, was könnte diese Summe $\$X$ sein? Zum Beispiel, wenn Sie die Wahl haben zwischen einer Wette bei der Sie bei gleicher Wahrscheinlichkeit $\$100$ gewinnen oder $\$50$ verlieren könnten und Ihnen keine Wette gleichgültig ist, dann entspricht Ihr X $\$100$. Geben Sie den Wert von X im unten stehenden Dialogfeld 'Positives Ergebnis' ein. Beachten Sie, dass je

größer X ist, desto weniger risikoscheu sind Sie, während ein kleineres X darauf hindeutet, dass Ihre Scheu vor dem Risiko höher ist.

Geben Sie die erforderlichen Inputs ein, wählen Sie den $U(x)$ -Typ aus und klicken Sie auf 'Nutzen berechnen', um die Ergebnisse zu erhalten. Sie können auch die errechneten $U(x)$ -Werte auf dem Entscheidungsbaum anwenden, um ihn erneut auszuführen, oder den Baum zurück zur Verwendung der Payoffs-Erwartungswerte schalten.

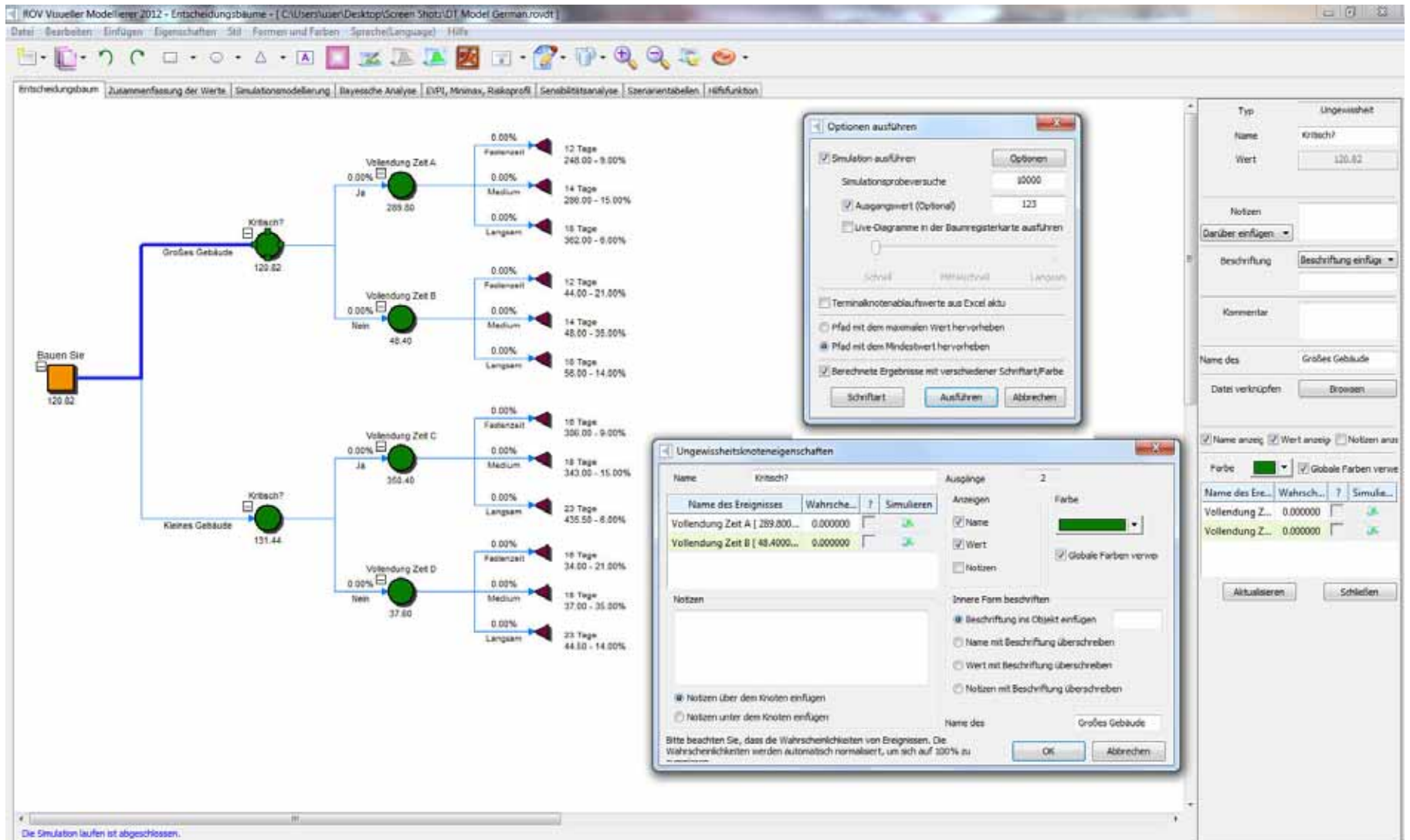


Bild 5.61 – ROV Entscheidungsbaum (Entscheidungsbaum)

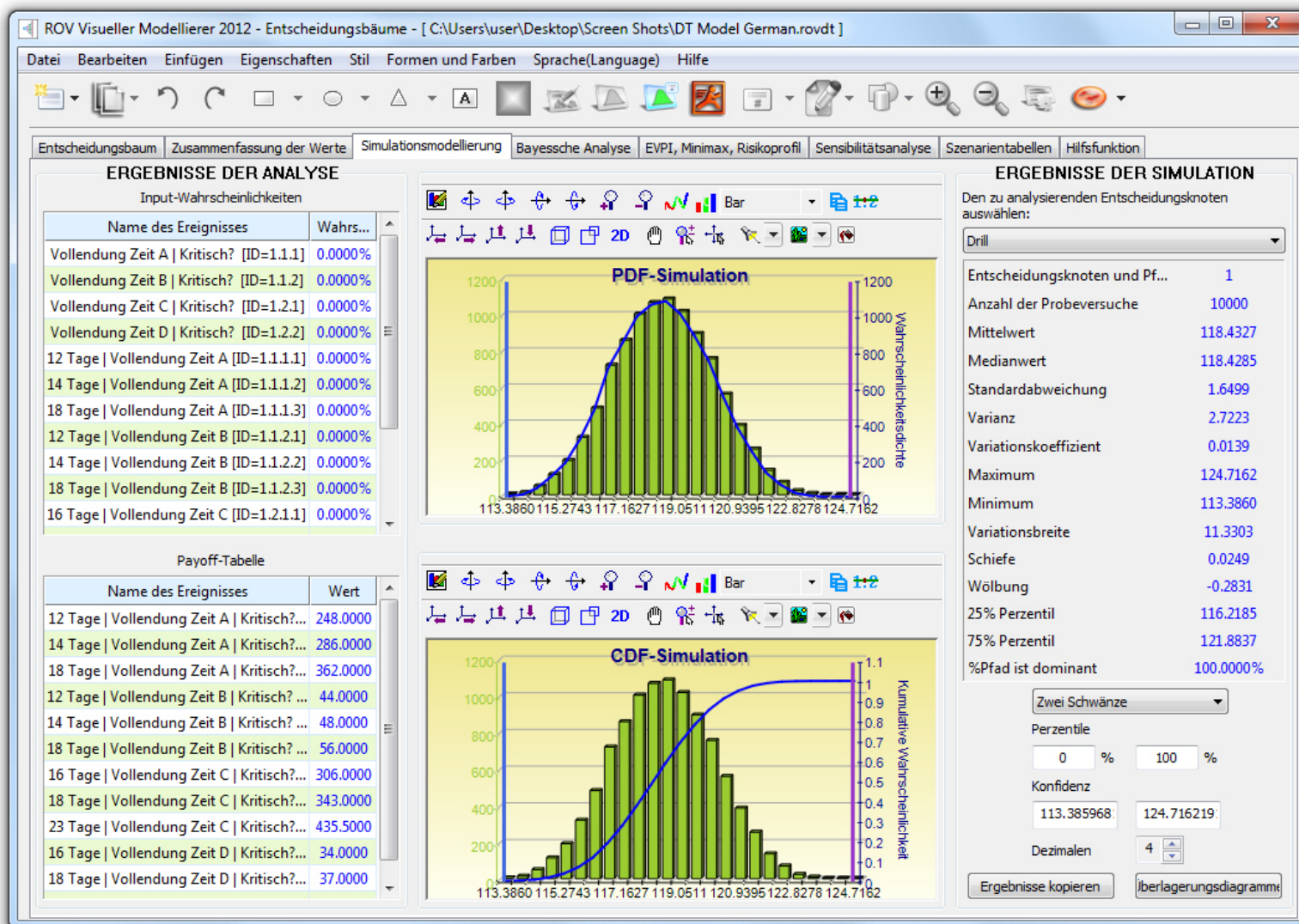


Bild 5.62 – ROV Entscheidungsbaum (Simulationsergebnisse)

ROV Visueller Modellierer 2012 - Entscheidungsbäume - [C:\Users\user\Desktop\Screen Shots\DT Model.rovdt]

Datei Bearbeiten Einfügen Eigenschaften Stil Formen und Farben Sprache(Language) Hilfe

Entscheidungsbaum Zusammenfassung der Werte Simulationsmodellierung Bayessche Analyse EVPI, Minimax, Risikoprofil Sensibilitätsanalyse Szenariotabellen Hilfsfunktion

Dieses bayessches Analysetool kann auf zwei beliebigen Ungewissheitsereignissen, die entlang einem Pfad verknüpft sind, ausgeführt werden. Zum Beispiel: im Beispiel rechts sind die Ungewissheiten A und B verknüpft, wobei Ereignis A als Erstes in der Zeitlinie und Ereignis B als Zweites eintritt. Das erste Ereignis A ist Marktforschung mit zwei Ausgängen (erfolgreich oder nicht erfolgreich). Das zweite Ereignis B ist Marktbedingungen, ebenfalls mit zwei Ausgängen (stark und schwach). Dieses Tool wird verwendet, um die gemeinsamen, die marginalen und die bayesschen a-posteriori aktualisierten Wahrscheinlichkeiten durch die Eingabe von vorherigen Wahrscheinlichkeiten und zuverlässigkeitsbedingten Wahrscheinlichkeiten zu berechnen. Man kann aber auch Zuverlässigkeitswahrscheinlichkeiten berechnen, wenn man a-posteriori aktualisierte bedingte Wahrscheinlichkeiten besitzt. Wählen Sie die relevante gewünschte Analyse weiter unten aus und klicken Sie auf 'Beispiel laden'. So können Sie Folgendes visualisieren: die Musterinputs, die der ausgewählten Analyse entsprechen, die im Raster rechts angezeigten Ergebnisse und welche Ergebnisse als Inputs in dem Entscheidungsbaum in der Abbildung verwendet sind.

☒ Bayessche a-posteriori Wahrscheinlichkeiten bei angegebenen gemeinsamen Vor- und Zuverlässigkeitswahrscheinlichkeiten berechnen (üblicher)

☐ Gemeinsame Zuverlässigkeitswahrscheinlichkeiten bei angegebenen Vor- und a-posteriori-Wahrscheinlichkeiten berechnen (seltener)

SCHRITT 1: Den Namen des ersten und den Namen des zweiten Ungewissheitsereignisses eingeben und die Anzahl der Wahrscheinlichkeitsereignisse (Naturzustände oder Ausgänge) für jedes Ereignis auswählen.

Name des ersten Ereignisses: Wahrscheinlichkeitsereignisse oder

Name des zweiten Ereignisses: Wahrscheinlichkeitsereignisse oder

SCHRITT 2: Die Namen aller Wahrscheinlichkeitsereignisse oder Ausgänge eingeben.

Zustände	Market Research	Market Conditions
1	Favorable	Strong
2	Unfavorable	Weak

SCHRITT 3: Die Vorwahrscheinlichkeiten und die bedingten Wahrscheinlichkeiten für jedes Ereignis oder Ausgang des zweiten Ereignisses eingeben. Die Wahrscheinlichkeiten müssen sich auf 100% summieren.

Bedingte Wahrscheinlichkeiten (Zuverlässigkeiten)

Ereignisse	Vor-W(x)	Favorable	Unfavorable	SUMME
Strong	45.00%	60.00%	40.00%	100.00%
Weak	55.00%	30.00%	70.00%	100.00%
SUMME	100.00%			

Gespeichertes Modell

Name:

Beispiel laden

Berechnen

Decision

First Uncertainty

Market Research

Second Uncertainty (F)

Second Uncertainty (U)

Favorable

Unfavorable

Strong Market

Weak Market

Strong Market

Weak Market

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%

Ergebnisse der bayesschen Analyse

Vorwahrscheinlichkeiten und zuverlässigkeitsbedingte

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

Gemeinsame und marginale Wahrscheinlichkeiten

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

A-posteriori oder aktualisierte Wahrscheinlichkeiten

Wahr. (Strong Favorable)	62.07%
Wahr. (Weak Favorable)	37.93%
Wahr. (Strong Unfavorable)	31.86%
Wahr. (Weak Unfavorable)	68.14%

Bild 5.63 – ROV Entscheidungsbaum (Bayessche Analyse)

ROV Visueller Modellierer 2012 - Entscheidungsbäume - [C:\Users\user\Desktop\Screen Shots\DT Model.rovdt]

Datei Bearbeiten Einfügen Eigenschaften Stil Formen und Farben Sprache(Language) Hilfe

Entscheidungsbaum Zusammenfassung der Werte Simulationsmodellierung Bayessche Analyse **EWPI, Minimax, Risikoprofil** Sensibilitätsanalyse Szenarientabellen Hilfsfunktion

Erwartungswert der perfekten Information, Minimax und Maximin Analyse, Risikoprofile und Wert der unvollständigen Information

Dieses Tool berechnet den Erwartungswert der perfekten Information (EWPI), die Minimax- und Maximin-Analyse und auch das Risikoprofil und den Wert der unvollständigen Information. Dabei geben Sie die Anzahl der Entscheidungsverzweigungen oder -strategien unter Berücksichtigung (z.B. der Bau einer großen, mittleren oder kleinen Anlage) und die Anzahl der ungewissen Ereignisse oder Zustände der Naturausgänge (z.B. guter Markt, schlechter Markt) und auch die erwarteten Payoffs unter jedes Szenario ein.

Inputhypothesen

Entscheidungsverzweigungen 3

Ungewissheitsereignisse oder -zustände 2

Wahrscheinl...	Zustand 1	Zustand 2	SUMME
	80%	20%	100%

Beispiel laden

Payoffs	Zustand 1	Zustand 2	Mittelwert
Entscheidun...	8	7	7.80
Entscheidun...	14	5	12.20
Entscheidun...	20	-9	14.20
Maximum	20.00	7.00	

Berechnen

Minimax- und Maximin-Analyse

Minimax (Minimierung des maximalen Regrets) und Maximin (Maximierung des minimalen Payoffs) sind zwei alternative Ansätze zur Ermittlung des optimalen Entscheidungspfads. Diese beiden Ansätze werden nicht oft verwendet, liefern jedoch zusätzliche Erkenntnisse über den Entscheidungsfindungsprozess. Geben Sie die Anzahl der existierenden Entscheidungsverzweigungen oder -pfade (z.B. der Bau einer großen, mittleren oder kleinen Anlage) und die Ungewissheitsereignisse oder Naturzustände unter jeden Pfad (z.B. gute Wirtschaft gegen schlechte Wirtschaft) ein. Dann vervollständigen Sie die Payoff-Tabelle für die verschiedenen Szenarien und berechnen Sie die Minimax- und Maximin-Ergebnisse. Sie können auch auf Beispiel laden klicken, um eine Beispielsberechnung zu sehen.

Payoffs	Zustand 1	Zustand 2	Minimum
Entschei...	8	7	7.00
Entschei...	14	5	5.00
Entschei...	20	-9	-9.00

Regret	Zustand 1	Zustand 2	Maximum
Entschei...	12.00	0.00	12.00
Entschei...	6.00	2.00	6.00
Entschei...	0.00	16.00	16.00

MINIMAX	6.00	Pfad 2 ist optimal
MAXIMIN	7.00	Pfad 1 ist optimal

Risikoprofil

Strategie 1 Risikoprofil

Payoff	Wahrscheinlic...
248.00	9.00%
286.00	15.00%
362.00	6.00%
44.00	21.00%
48.00	35.00%
56.00	14.00%
Summe der Wahrscheinlich...	100.00%
Erwartungswert	120.82

Strategie 2 Risikoprofil

Payoff	Wahrscheinlic...
306.00	9.00%
343.00	15.00%

Erwartungswert der unvollständigen Information 10.62

Gespeichertes Modell

Name

HINZUFÜGEN

DEL

Erwartungswert der perfekten Information von Naturzuständen 17.40

Erwartungswert ohne die perfekte Information von Naturzuständen 14.20

Erwartungswert der perfekten Information 3.20

Bild 5.64 – ROV Entscheidungsbaum (EWPI, MINIMAX, Risikoprofil)

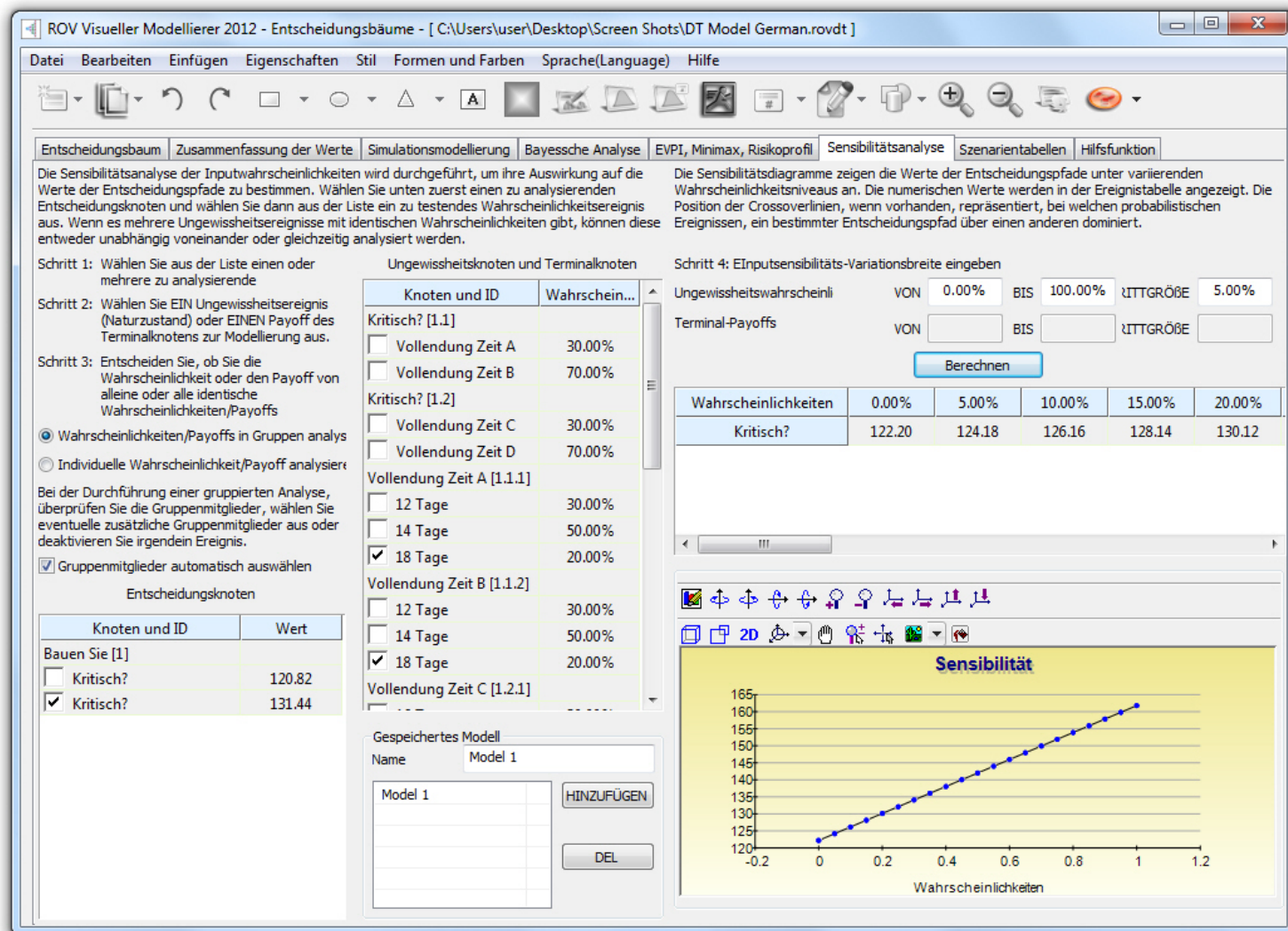


Bild 5.65 – ROV Entscheidungsbaum (Sensibilitätsanalyse)

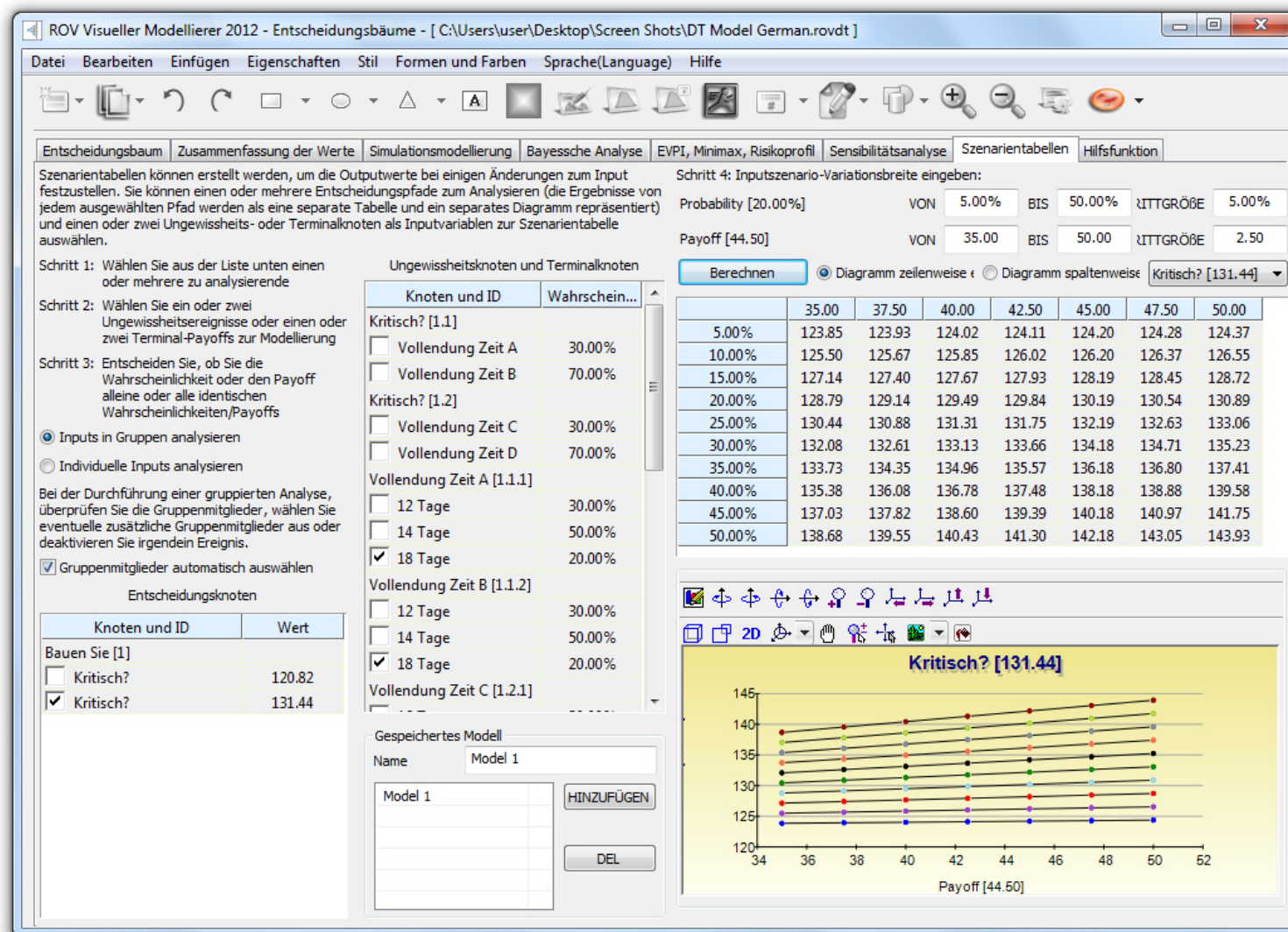


Bild 5.66 – ROV Entscheidungsbaum (Szenarietabellen)

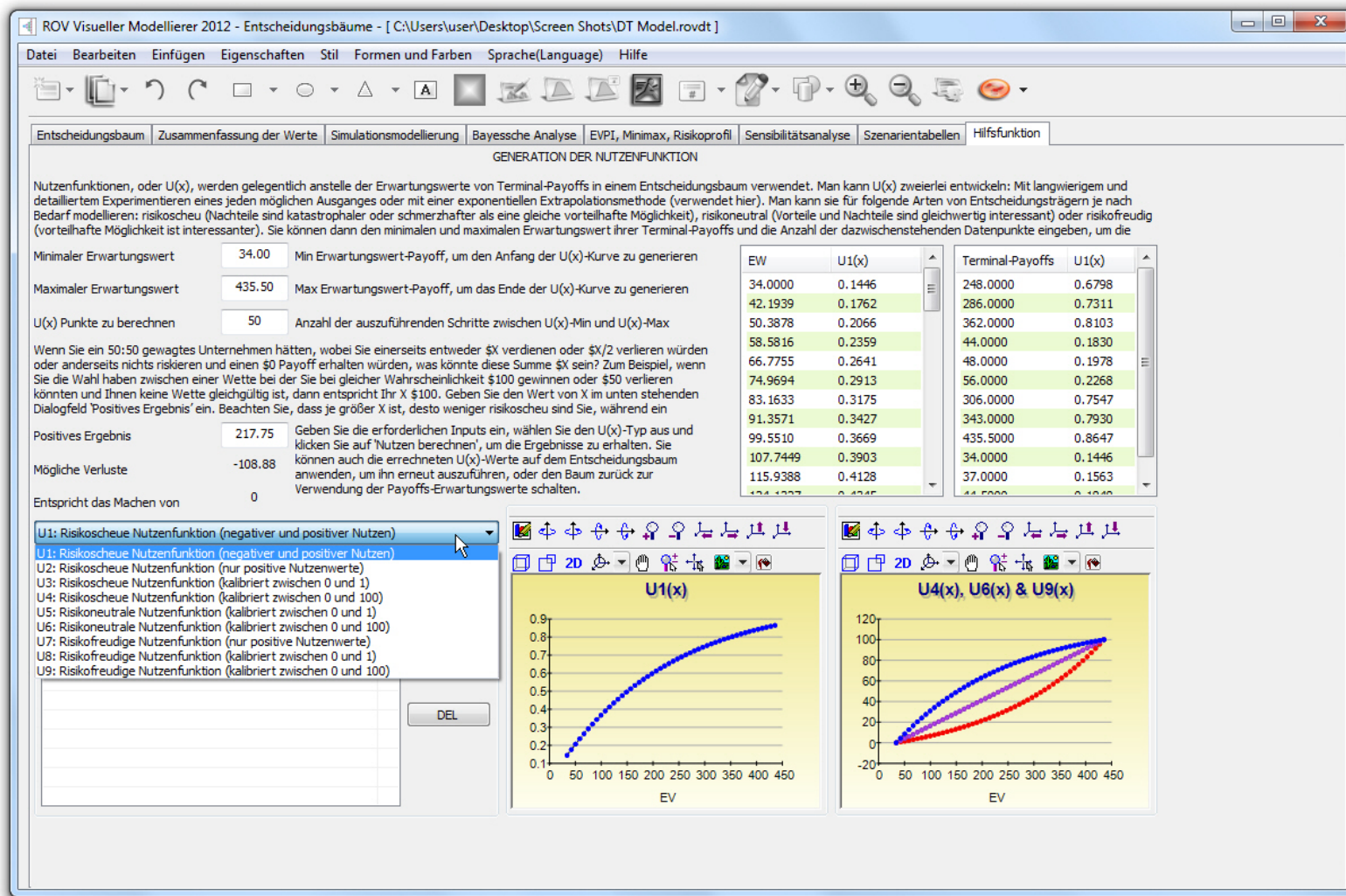


Bild 5.67 – ROV Entscheidungsbaum (Nutzenfunktionen)

HILFBREICHE TIPPS UND TECHNIKEN

Folgend sind einige schnelle hilfreiche Tipps und Verknüpfungstechniken für fortgeschrittene Risiko Simulator User. Für Details zur Anwendung spezifischer Tools, bitte die relevanten Abschnitte des User Manuals nachschlagen.

TIPPS: Annahmen (Festgesetzte Input- Annahme Benutzeroberfläche)

- Schneller Sprung—eine Verteilung auswählen und einen beliebigen Buchstaben eintippen—es wird dann an die erste Verteilung die mit diesem Buchstaben anfängt springen (z.B. Normal anklicken und W eintippen, und es wird Sie zur Weibull Verteilung führen).
- Rechtsklick Ansichten—beliebige Verteilung auswählen, rechtsklicken und die verschiedene Ansichten der Verteilung auswählen (große Icons, kleine Icons, Liste).
- Tabulator um Charts zu aktualisieren—nachdem Sie einige neue Input Parameter eingegeben haben (z.B. Sie tippen einen neuen Mittelwert oder Standardabweichungswert ein), Tab auf der Tastatur drücken oder irgendwo auf dem User Interface entfernt von dem Input Eingabefeld klicken um die automatische Aktualisierung des Verteilungscharts zu sehen.
- Korrelationen Eingeben—Sie können hier sofort Korrelationen paarweise eingeben (die Säulen sind je nach bedarf größenveränderbar), das multipel Verteilungsanpassungs-Tool anwenden um alle Korrelationen paarweise automatisch zu berechnen und einzugeben, oder nach der Einstellung einiger Annahmen das Edit- Korrelation-Tool anwenden um Ihre Korrelation Matrix einzugeben.
- Gleichstellungen in eine Annahme-Speicherzelle—lediglich leere Zellen oder Zellen mit statischen Werten können als Annahmen eingestellt werden; jedoch könnte es vorkommen, dass eine Funktion oder Gleichstellung in einer Annahmezelle benötigt wird, und dies lässt sich machen wenn Sie zuerst die Inputannahme in die Zelle eingeben und dann die Gleichstellung oder Funktion eintippen (wenn die Simulation am Laufen ist, die simulierte Werte werden die Funktion ersetzen, und nachdem die Simulation vervollständigt ist, die Funktion oder Gleichstellung wird wieder aufgezeigt).

TIPPS: Kopieren und Einfügen TIPS: Copy and Paste

- Kopieren und Einfügen mittels *Escape Abbruch*—wenn Sie eine Zelle auswählen und die Risiko Simulator Kopier Funktion anwenden, wird alles in die Zwischenablage kopiert, inklusive Zellwert, Gleichstellung, Funktion, Farbe, Font und Größe, sowie Risiko Simulator Annahmen, Prognosen oder

Entscheidungsvariablen. Dann, so wie Sie die Risiko Simulator Einfügen Funktion anwenden, haben Sie zwei Alternativen. Die erste Alternative ist die Risiko Simulator Einfügen Funktion direkt anzuwenden, und alle Zellwerte, Farbe, Font, Gleichstellung, Funktionen und Parameter werden in die neue Zelle eingefügt. Die zweite Alternative ist erst Escape **Abbruch** auf der Tastatur anzuklicken, und dann Risiko Simulator Einfügen anzuwenden. **Abbruch** teilt Risiko Simulator mit, dass Sie nur die Risiko Simulator Annahme, Prognose oder Entscheidungsvariable einfügen möchten, und nicht die Zellwerte, Farbe, Gleichstellung, Funktion, Font usw. **Abbruch** vor dem Einfügen zu drücken lässt zu, dass Sie die Werte und Berechnungen der Zielzelle beibehalten, und fügt nur die Risiko Simulator Parameter ein.

- Kopieren und Einfügen auf Multiple Zellen—Sie können multiple Zellen für kopieren und einfügen auswählen (mit angrenzenden und nicht angrenzenden Annahmen).

TIPPS: Korrelationen

- Annahme Einstellung (oder Festsetzung) (oder Festgesetzte Annahme)—Paarweise Korrelationen mit der Anwendung des Annahme-Input Einstellungsdialogs „Set Input Assumption Dialogs“ einstellen (ideal nur für die Eingabe mehrerer Korrelationen).
- Korrelationen Aufbereiten—Eine Korrelationsmatrix manuell einstellen durch eingeben oder einfügen aus der Windows Zwischenablage (ideal für große Korrelationen Matrizen und multiple Korrelationen).
- Multiple Verteilungsanpassung—Paarweise Korrelationen werden automatisch berechnet und eingegeben (ideal wenn multiple Variable Anpassung ausgeführt wird, Korrelationen automatisch berechnet werden, und um zu entscheiden was eine statistisch signifikante Korrelation ausmacht).

TIPPS: Datendiagnostik und Statistische Analyse

- Stochastische Parametereinschätzung—In den statistischen Analyse und Datendiagnostik Berichten ist ein Tab über stochastische Parametereinschätzungen der die Unbeständigkeit, Drift (Abweichung), Durchschnitts-Reversionsrate und Sprungdiffusions-Raten einschätzt, basierend auf historischen Daten. Beachten Sie, dass diese Parameter Ergebnisse ausschließlich auf die angewendeten historischen Daten bezogen sind, dass die Parameter sich mit der Zeit ändern könnten, und sind auch von der Menge der angepassten historischen Daten abhängig. Ferner, die Analyse Ergebnisse zeigen alle Parameter und implizieren nicht welches

stochastisches Prozess-Modell (z.B. Brownsche Bewegung, Durchschnitts-Reversion, Sprungdiffusion, oder gemischte Prozesse) das best passende ist. Der User muss je nach der Zeitreihen-Variable die vorausberechnete werden soll diesen Entschluss machen. Die Analyse kann nicht entscheiden welchen Prozess am besten ist, sondern dies kann nur der User (z.B. Das Brownsche Bewegung Prozess ist am besten um die Aktienpreise zu modellieren, aber die Analyse kann nicht determinieren ob die analysierte historische Daten von einer Aktie oder von einer anderen Variable ist - nur der User wird dies wissen). Zum Abschluss ein guter Hinweis: wenn sich eine bestimmte Parameter außerhalb des normalen Bereichs befindet, ist der Prozess der diese Input Parameter benötigt, wahrscheinlich nicht der richtige Prozess (z.B. wenn die Durchschnitts-Reversionsrate 110% beträgt, ist dies aller Wahrscheinlichkeit nach nicht der richtige Prozess, usw).

TIPPS: Verteilungsanalyse, Charts und Wahrscheinlichkeitstabellen

- Verteilungsanalyse—Angewendet um die PDF, CDF, und ICDF der 42 Wahrscheinlichkeitsverteilungen verfügbar im Risiko Simulator schnell zu berechnen, und eine Tabelle dieser Werte wiederzugeben.
- Verteilungscharts und Tabellen—Angewendet um *verschiedene Parameter der gleichen Verteilung* zu vergleichen (z.B., die Formen und PDF, CDF, ICDF Werte einer Weibull Verteilung mit Alpha und Beta von [2, 2], [3, 5] und [3.5, 8], und überlagern sie übereinander).
- Overlay Charts—angewendet um *verschiedene Verteilungen* zu vergleichen (theoretische Inputannahmen und empirisch simulierte Outputvoraussagen) und sie für einen visuellen Vergleich übereinander zu legen.

TIPPS: Effiziente Grenze

- Effiziente Grenzvariablen—Um die Grenzvariablen abzurufen, erst die Grenzen des Modells einstellen vor der Einstellung der effizienten Grenzvariablen.

TIPPS: Vorassagezellen

- Voraussagezellen mit keinen Gleichungen—Sie können Outputvoraussagen in Zellen ohne jegliche Gleichungen oder Werte einstellen (einfach die Warnmeldung ignorieren) aber beachten Sie, dass das resultierende Voraussagechart leer sein wird. Outputvoraussagen werden in leere Zellen typischerweise eingestellt wenn Macros da sind die berechnet werden, und die Zelle kontinuierlich aktualisiert wird.

TIPPS: Voraussageschart

- *Tab* gegen *Leertaste*—*Tab* auf der Tastatur drücken um das Voraussageschart zu aktualisieren und die Perzentil und Vertrauenswerte zu erhalten nachdem Sie einige Inputs eingegeben haben, und die *Leertaste* drücken um zwischen den verschiedenen Tabs in dem Voraussageschart zu rotieren.
- *Normal* gegen *Globale Ansicht*—Diese Ansichten anklicken um zwischen eine angesteuerte Schnittstelle und eine globale Schnittstelle zu rotieren wo alle Elemente der Voraussagecharts gleichzeitig sichtbar sind.
- *Kopieren*—Dies wird das Voraussageschart oder die gesamte globale Ansicht kopieren, je nachdem ob Sie sich in der normalen oder globalen Ansicht befinden.

TIPPS: Voraussagen

- *Zellen Link-Adresse*—Wenn Sie zuerst die Daten in der Tabellenkalkulation auswählen und dann ein Voraussage-Tool ausführen, wird die Zelladresse der ausgewählte Daten automatisch in die User Schnittstelle eingefügt, anderenfalls müssten Sie manuell in die Zelladresse eingeben oder das *Link* Icon anwenden um sich mit der relevanten Daten Speicherstelle zu verlinken.
- *Voraussage RMSE*—Dies als das allgemeingültige Fehlermaß bei multiple Voraussagemodelle für direkte Vergleiche der Genauigkeit jedes Modells.

TIPPS: Voraussagen: ARIMA

- *Voraussage Perioden*—die Zahl der exogenen Datenzeilen muss die Zeitreihen Datenzeilen um mindestens die erwünschten Voraussage Perioden übertreffen (z.B. wenn Sie 5 Perioden in die Zukunft voraussagen möchten und haben 100 Zeitreihen Datenpunkte, werden Sie mindestens 105 oder mehr Datenpunkte auf der exogenen Variable haben müssen), andernfalls einfach ARIMA ohne die exogene Variable ausführen um so viele Perioden wie Sie möchten vorauszusagen ohne jegliche Einschränkungen.

TIPPS: Vorassagen: Basic Ökonometrie

- *Variable Trennung mit Semikola*—Unabhängige Variablen trennen mittels eines Semikolons.

TIPPS: Vorassagen: Logit, Probit und Tobit

- Datenanforderungen—Die abhängige Variablen um Logit und Probit Modelle auszuführen darf ausschließlich binär sein (0 und 1), während das Tobit Modell kann binär und andere numerische Dezimalwerte aufnehmen. Die unabhängige Variablen für alle drei Modelle können beliebige numerische Werte aufnehmen.

TIPPS: Voraussagen: Stochastische Prozesse

- Vorgegebene Sample Inputs—Im Zweifelsfall, die vorgegebene Inputs als Anfangspunkt anwenden um Ihr eigenes Modell zu entwickeln.
- Statistisches Analyse Tool—Diese Tool anwenden um die Input Parameter in die stochastischen Prozessmodellen zu kalibrieren, indem Sie sie aus Ihren Rohdaten abschätzen.
- Stochastisches Prozessmodell—manchmal, wenn die stochastische Prozess Benutzerschnittstelle lange Zeit hängt, sind aller Wahrscheinlichkeit nach Ihre Inputs inkorrekt, und das Modell ist nicht richtig spezifiziert (z.B. wenn die Durchschnitts-Reversionsrate 110% beträgt, ist dies aller Wahrscheinlichkeit nach nicht der richtige Prozess, usw). Bitte versuchen Sie es mit anderen Inputs oder wenden Sie ein anderes Modell an.

TIPPS: Voraussagen: Trendlinien

- Vorassagen Ergebnisse—Um die vorausgesagten Werte zu sehen, scrollen Sie bis zum Ende des Berichts.

TIPPS: Funktionsaufrufe

- RS Funktionen—Es gibt Funktionen – „Input-Annahme einstellen“ und „Voraussage-Statistiken erhalten“, die Sie in Ihrer Excel Tabellenkalkulation anwenden können. Um diese Funktionen anzuwenden, müssen Sie zuerst RS Funktionen einstellen (Starten, Programme, Real Optionen Bewertung, Risiko Simulator, Tools und Funktionen Installieren) und dann, bevor Sie die RS Funktionen innerhalb Excel einstellen, eine Simulation ausführen. Für Beispiele wie diese Funktionen anzuwenden sind, auf dem Beispiel Modell 24 verweisen.

TIPPS: Einstiegsübungen und Einstiegsvideos

- Einstiegs-Übungen—Multiple Schritt-für-Schritt praktische Beispiele und Ergebnis-Interpretations-Übungen sind in dem Starten, Programme, Real Optionen Bewertung, Risiko Simulator Tastaturkürzel verfügbar. Diese Übungen sind dafür gedacht, Sie schnell mit der Software Anwendung auf Hochtouren bringen.
- Einstiegs-Videos—Diese sind alle kostenlos auf unserer Website über www.realoptionsvaluation.com/download.html oder www.rovdownloads.com/download.html verfügbar.

TIPPS: Hardware ID

- HWID Kopieren Rechts Anklicken—In die „Lizenz Installieren“ Benutzeroberfläche, die HWID auswählen oder doppelklicken um deren Wert zu selektieren, rechtsklicken um zu kopieren, oder die E-Mail HWID anklicken um eine E-Mail mit der HWID zu generieren.
- Problemlöser—Problemlöser aus dem Starten, Programme, Real Optionen Bewertung, Risiko Simulator Ordner ausführen, und das Get HWID/HWID holen. Tool laufen lassen um die HWID Ihres Computers zu bekommen.

TIPPS: Latin Hypercube Sampling (LHS) gegen Monte Carlo Simulation (MCS)

- Korrelationen—Wenn paarweise Korrelationen zwischen Input Annahmen eingestellt werden sollen, empfehlen wir die Anwendung der Monte Carlo Einstellung im Risiko Simulator, Optionen Menü. Latin Hypercube Sampling ist nicht mit der korrelierten Kopula Methode für Simulation kompatibel. *Correlations—when setting pairwise correlations among input assumptions, we recommend using the Monte Carlo setting in the Risk Simulator, Options menu. Latin Hypercube Sampling is not compatible with the correlated copula method for simulation.*
- LHS Behälter (Kästen) —Eine höhere Zahl der Behälter wird die Simulation verlangsamen wobei ein einheitlicheren Satz der Simulationsergebnissen gewährleistet wird.
- Beliebigkeit—alle beliebige Simulationstechniken im Optionen Menü wurden getestet - alle sind gute Simulatoren und nähern das gleiche Niveau der Beliebigkeit an bei der Ausführung einer größeren Anzahl der Proben.

TIPPS: Online Ressourcen TIPS: Online Resources

- Bücher, Einstiegs-Videos, Modelle, kommerzielle IT-Publikationen—
kostenlos auf unserer Website über
www.realoptionsvaluation.com/download.html oder
www.rovdownloads.com/download.html.

VERZEICHNIS

- abhängige Variable*, 167, 168, 170
- Aktienpreis*, 171, 172
- Aktiva*, 134, 135
- Aktivaklassen*, 134, 135
- Aktivum*, 135, 136
- Alpha*, 167
- Analyse*, 135, 138, 168, 170, 175
- Anderson-Darling-Test*, 157
- anhalten*, 28
- Annahmen*, 171
- annualisierten*, 135
- ARIMA*, 8, 85, 100, 101, 102, 103, 105, 106, 107, 109, 111, 113, 116, 118, 169
- Aufteilung*, 134, 135, 136
- Ausführung*, 168, 170
- Ausreißer*, 166, 167, 168, 169, 170
- Autokorrelation*, 169, 173
- bedeuten*, 170
- Bereich*, 26, 44, 58, 66, 121, 124, 135, 138, 141, 168
- Bericht*, 23, 88, 92, 95, 98, 102, 144, 152, 155, 164, 171
- Bestimmungskoeffizient*, 166
- Beta*, 58
- Bevölkerung*, 167, 171
- Binomial*, 51, 52, 53, 54, 55
- Bootstrap*, 8, 158, 160, 161
- Box-Jenkins*, 8, 100, 106
- Brownsche Bewegung*, 171, 172
- Chi-Quadrat-Test*, 157
- Crystal Ball*, 50, 141, 156, 157
- Daten*, 31, 36, 37, 39, 48, 49, 58, 70, 77, 80, 81, 85, 86, 87, 88, 90, 91, 92, 95, 98, 99, 100, 101, 102, 106, 107, 153, 154, 155, 157, 158, 160, 163
- Delphi*, 80, 153
- Delphi-Methode*, 153
- diskret*, 8, 49, 120
- diskrete*, 51, 53, 122, 128, 157, 172
- Diskrete*, 49
- Dispersion*, 38, 43
- Dreieck*, 19, 50, 62
- drittes Moment*, 43, 45
- Einschränkungen*, 136
- Einzel*, 168
- Einzel-*, 138
- Einzel-Aktivum SLS*, 8
- einzelne*, 175
- e-mail*, 2
- E-Mail*, 10
- Entscheidungen*, 138
- entscheidungsvariable*, 122
- Entscheidungsvariable*, 120, 121, 122, 123, 125, 128, 129
- Entscheidungsvariablen*, 135
- Erlang*, 66, 67
- erster Moment*, 44
- erstes Moment*, 43
- Excel*, 8, 9, 10, 21, 22, 23, 31, 36, 37, 85, 87, 92, 98, 102, 106, 107, 109, 111, 113, 116, 118, 120, 143
- Extrapolation*, 8, 98, 99
- Exzesswölbung*, 46, 51, 52, 53, 55, 56, 57, 59, 61, 62, 64, 65, 66, 69, 70, 71, 72, 76, 77, 78
- Fehler*, 8, 9, 22, 23, 28, 31, 40, 72, 88, 90, 98, 101, 103, 158, 166, 167, 170, 171, 173
- Fisher-Snedecor*, 65
- Flexibilität*, 138
- Fluktuationen*, 172
- Fluktuationen*, 166
- Frequenz*, 48, 49
- Funktionen*, 169
- Funktionen von*, 169
- Galerie*, 25, 26
- Gamma*, 59, 60, 61, 66, 67, 76, 78, 88
- Ganzzahl*, 23, 52, 55, 61, 67, 77, 88, 120, 122
- ganzzahlig*, 8
- geometrisch*, 123
- geometrische*, 53, 55
- Geometrische*, 53
- geometrischen*, 53, 70
- geometrisches Mittel*, 135
- Gleichung*, 168, 171, 173
- Gültigkeit von*, 169
- Güte-der-Anpassung*, 170
- Güte-der-Anpassung*, 169
- Güte-der-Anpassung Tests*, 169
- Heteroskedastizität*, 166, 167, 169, 170
- Histogramm*, 48, 49
- Holt-Winter*, 88, 89
- hypergeometrische*, 54
- Hypergeometrische*, 54
- hypergeometrischen*, 54

Hypothese, 8, 19, 23, 24, 25, 26, 27, 28, 29, 50, 60, 65, 76, 92, 101, 103, 121, 136, 150, 152, 155, 157, 158, 161, 162
Hypothesen, 134, 167
Ikone, 10, 24, 26, 27, 28, 125, 129
Ikonen, 136
Inflation, 169, 172
Inputs, 136
Installation, 9, 10
Investitionen, 134
Ja/Nein, 51
Kausalität, 174
kleinste Quadrate, 167
Kolmogorov-Smirnov-Test, 157
Konfidenzintervall, 32, 34, 40, 76, 158, 160
Kontinuierlich, 49
Korrelation, 19, 23, 36, 37, 38, 50, 101, 103, 152, 153, 169, 173, 174
Korrelationen, 174
Korrelationskoeffizient, 174
linear, 168
lineare, 167, 170
lineares, 173, 174
Ljung-Box Q- Statistiken, 169
logistische, 68, 69
Lognormal, 69, 70
Management, 138
Markt, 168, 172, 174
Matrix, 173
mehrfach, 134, 138
mehrfache, 173, 175
mehrfache Regression, 173
mehrfache Variablen, 175
Methode, 135, 136, 140, 168, 173
Mittelwert, 69, 70, 71, 167, 171
Mix, 173
Modell, 135, 136, 142, 166, 168, 169, 170
Modelle, 168
modellieren, 171
Monte Carlo, 19, 51
Monte-Carlo, 19, 40, 50
Multikollinearität, 166, 173
Multinomial SLS, 9
multivariate, 90, 91, 92, 98, 101, 102
Mun, 1, 2, 8, 88, 91, 92, 95, 102
N, 95
negative Binomial, 55
nicht linear, 168, 174
nicht lineare, 167
normal, 60
Normal, 19, 26, 36, 41, 46, 50, 52, 60, 69, 70, 71, 77, 155, 158, 167
Nullhypothese, 167, 169, 171
obere, 135
optimal, 168
optimale, 138
optimale Entscheidung, 138
Optimierung, 8, 20, 120, 121, 122, 123, 124, 125, 127, 128, 129, 132, 134, 135, 136, 138, 157
Optimierungen, 122
Option, 8, 31, 88, 89, 147, 155, 163
Parameter, 70, 172
Pareto, 72
Pearson, 36, 37
Poisson, 57, 63, 66
Portfolio, 134, 138
Präzision, 8, 22, 28, 31, 40
Preis, 94
Probeversuche, 19, 22, 23, 28, 29, 40, 50, 51, 52, 53, 54, 55, 64, 120, 121, 136, 158
profil, 129
Profil, 21, 22, 23, 24, 37, 88, 125, 136, 155
Punktschätzung, 138
p-Wert, 174
p-Werte, 169
Querschnitt, 81, 98
Rangkorrelation, 174
Rate, 169, 172
Regression, 8, 90, 91, 92, 98, 101, 102
Regression der kleinsten Quadrate, 167
Regressionsanalyse, 166, 167, 168
relative Renditen, 135
Reliability, 141
Rendite, 134, 135, 136
Renditen, 134, 135, 136, 168
Risiko, 134, 135, 136
Risiko Simulator, 136
Rückkehr zum Mittelwert, 95
Saisonalität, 170
Schätzungen, 167, 168, 169
Schiefe, 43, 45, 46, 51, 52, 53, 55, 56, 57, 59, 60, 61, 62, 64, 65, 66, 69, 70, 71, 72, 76, 77, 78, 160
Sensibilität, 8, 141, 145, 150, 151, 152, 153
Signifikanz, 167, 169, 171, 174
Simulation, 8, 19, 20, 21, 22, 23, 24, 28, 29, 30, 35, 36, 37, 38, 39, 40, 48, 50, 51, 81, 88, 95, 120, 121, 122, 124, 125, 129, 132, 134, 136, 138, 141, 145, 150, 151, 152, 153, 155,

158, 160, 161, 163, 164, 166, 174, 175, 180,
 183, 185
 SLS, 8
 Spearman, 36, 37
 speichern, 10, 22, 163
 Spezifikationsfehler, 166
 Spinnennetz, 8, 143, 144, 147, 150
 Sprung-Diffusion, 95
 Standardabweichung, 19, 31, 38, 41, 44, 45,
 46, 50, 52, 53, 54, 57, 60, 61, 62, 64, 65, 66,
 70, 71, 76, 95, 121, 155, 160, 171, 174
 Standardabweichungen, 162
 statisch, 171
 Statistiken, 29, 31, 38, 40, 43, 44, 102, 103,
 121, 155, 158, 160
 Stichprobe, 171
 Stichproben, 167
 stochastisch, 8, 94, 95, 120, 121, 122, 125,
 129, 132, 171, 172
 stochastische, 120, 121, 135, 136, 138, 140,
 166
 stochastische Optimierung, 135, 136, 138, 140
 Symbolleiste, 10, 24, 27, 28
 symmetrisch, 167
 Titel, 22
 Tornado, 8, 141, 143, 144, 145, 147, 150, 151,
 152, 153
 Trends, 171
 t-Statistiken, 173
 t-Verteilung, 76
 Typen von, 134, 171
 Übernahme, 168
 Umsätze, 169
 unabgängige Variable, 173
 unabhängige Variable, 167, 168, 170, 173
 uniform, 123
 Uniform, 19, 50, 53, 77, 135, 153
 untere, 135
 Varianz, 166, 167
 Verhalten, 171
 Verhältnis, 134, 136
 Verkauf, 170
 Verteilung, 19, 24, 25, 26, 27, 29, 36, 38, 43,
 45, 46, 48, 49, 50, 51, 52, 53, 54, 55, 57, 58,
 59, 60, 61, 62, 63, 64, 65, 66, 67, 68, 69, 70,
 71, 72, 76, 77, 78, 80, 95, 121, 136, 138, 152,
 153, 154, 155, 156, 157, 158, 160, 171
 Verteilungen, 134
 Verteilung, 55
 Verteilung, 65
 Verzögerungen, 169, 170
 vierter Moment, 46
 viertes Moment, 43
 Volatilität, 172
 Vorausberechnen, 172
 Vorausberechnung, 19, 23, 27, 28, 29, 30, 31,
 32, 35, 38, 40, 43, 50, 80, 81, 88, 98, 100, 102,
 106, 107, 121, 141, 151, 152, 158, 160, 162,
 163, 169
 Vorausberechnungen, 169
 Vorausberechnungsstatistiken, 158
 Vorausberechnungsstatistiken, 29, 121
 Vorausberechnung, 95, 166
 Vorhersage, 167, 168
 Wachstum, 134, 172
 Wachstumsrate, 172
 Wahrscheinlichkeit, 8, 19, 29, 32, 33, 34, 45,
 48, 49, 50, 51, 52, 53, 54, 55, 56, 58, 60, 65,
 70, 77
 Weibull, 77, 78
 Wert, 134, 135, 136, 167, 169, 171, 172, 173,
 174
 Werte, 134, 135, 167, 168, 169, 171, 174
 Wölbung, 46, 60
 Zeitreihe, 100
 Zeitreihen, 81, 86, 88, 95, 98, 99, 101, 102,
 169, 171, 172
 Zeitreihen, 8
 Zeitreihendaten, 169, 171, 172
 Zentrum der, 168
 Ziel, 136
 Zins, 169, 171, 172
 Zinssätze, 169, 171, 172
 Zufall, 171, 172
 Zufallszahl, 23, 50
 Zufallszahlen, 19
 zweiter Moment, 45
 zweites Moment, 43