

IBM SPSS Modeler 18.4 データベース内
マイニング・ガイド



注

本書および本書で紹介する製品をご使用になる前に、[93 ページの『特記事項』](#)に記載されている情報をお読みください。

<!--製品情報-->

本書は、IBM® SPSS® モデラー バージョン 18 リリース 4 モディフィケーション 0、および新しい版で明記されていない限り、以降のすべてのリリースおよびモディフィケーションに適用されます。

© Copyright International Business Machines Corporation .

目次

前書き.....	vii
第 1 章 IBM スポス モデラー について.....	1
IBM スポス モデラー 製品.....	1
IBM スポス モデラー.....	1
IBM スポスモデラー・サーバー.....	1
IBM スポス モデラー 管理コンソール.....	2
IBM スポス モデラー・バッチ.....	2
IBM スポスモデラー・ソリューション・パブリッシャ.....	2
IBM スポス コラボレーションおよびデプロイメント・サービス用の IBM スポスモデラー・サーバーアダプター.....	2
IBM スポス モデラー のエディション.....	2
資料.....	3
スポス モデラー・専門的 ドキュメント.....	3
スポスモデラー・保険料 ドキュメント.....	4
アプリケーションの例.....	4
Demos フォルダー.....	4
ライセンスの追跡.....	4
第 2 章データベース内マイニング.....	5
データベース・モデル作成の概要.....	5
必要事項.....	5
モデルの構築.....	6
データの準備.....	6
モデル・スコアリング.....	6
データベース・モデルのエクスポートと保存.....	7
モデルの整合性.....	7
生成された SQL の表示とエクスポート.....	7
第 3 章 Microsoft Analysis Services によるデータベース・モデリング.....	9
IBM スポス モデラー と Microsoft Analysis Services.....	9
Microsoft Analysis Services との統合の要件.....	10
Analysis Services との統合を有効にする.....	11
Analysis Services でのモデル構築.....	13
Analysis Services モデルの管理.....	13
すべてのアルゴリズム・ノードに共通の設定.....	15
MS ディジション・ツリーのエキスパート・オプション.....	15
MS クラスタリングのエキスパート・オプション.....	15
MS Naive Bayes のエキスパート・オプション.....	15
MS Linear Regression のエキスパート・オプション.....	16
MS Neural Network のエキスパート・オプション.....	16
MS Logistic Regression エキスパート・オプション.....	16
MS アソシエーション・ルール・ノード.....	16
MS Time Series ノード.....	16
MS Sequence Clustering ノード.....	18
Analysis Services モデルのスコアリング.....	19
すべての Analysis Services モデルに共通の設定.....	19
MS Time Series モデル・ナゲット.....	20
MS Sequence Clustering のモデル・ナゲット.....	21
モデルのエクスポートとノードの生成.....	21

Analysis Services マイニングの例.....	21
ストリームの例 :ディシジョン・ツリー.....	21
第 4 章 Oracle Data Mining によるデータベース・モデリング.....	25
Oracle Data Mining について.....	25
Oracle との統合の要件.....	25
Oracle との統合を有効にする.....	26
Oracle Data Mining を使用してモデルを構築する.....	27
Oracle モデルの「サーバー」オプション.....	28
誤分類コスト.....	28
Oracle Naive Bayes.....	29
Naive Bayes の「モデル」オプション.....	29
Naive Bayes の「エキスパート」オプション.....	29
Oracle Adaptive Bayes.....	29
Adaptive Bayes の「モデル」オプション.....	30
Adaptive Bayes の「エキスパート」オプション.....	30
Oracle Support Vector Machine (SVM).....	31
Oracle SVM の「モデル」オプション.....	31
Oracle SVM の「エキスパート」オプション.....	31
Oracle SVM の「重み」オプション.....	32
Oracle 一般化線型モデル (GLM)*.....	32
Oracle GLM の「モデル」オプション.....	33
Oracle GLM の「エキスパート」オプション.....	33
Oracle GLM の「重み」オプション.....	33
Oracle ディシジョン・ツリー.....	34
ディシジョン・ツリーのモデル・オプション.....	34
ディシジョン・ツリーのエキスパート・オプション.....	35
Oracle O-Cluster.....	35
O-Cluster の「モデル」オプション.....	35
O-Cluster の「エキスパート」オプション.....	35
Oracle K-Means.....	36
K-Means の「モデル」オプション.....	36
K-means の「エキスパート」オプション.....	36
Oracle 非負数マトリックス因数分解 (NMF).....	36
NMF の「モデル」オプション.....	37
NMF の「エキスパート」オプション.....	37
Oracle Apriori.....	37
Apriori の「フィールド」オプション.....	37
Apriori の「モデル」オプション.....	38
Oracle 最小記述長 (MDL).....	39
MDL のモデル・オプション.....	39
Oracle Attribute Importance (AI).....	39
AI のモデル・オプション.....	40
AI の選択オプション.....	40
AI モデル・ナゲットの「モデル」タブ.....	40
Oracle モデルの管理.....	40
Oracle モデル・ナゲットの「サーバー」タブ.....	40
Oracle モデル・ナゲットの「要約」タブ.....	41
Oracle モデル・ナゲットの「設定」タブ.....	41
Oracle モデルのリスト作成.....	41
Oracle Data Miner.....	42
データの準備.....	42
Oracle Data Mining の例.....	43
ストリームの例 :データのアップロード.....	43
ストリームの例 :データの調査.....	43
ストリームの例 :モデルの作成.....	44
ストリームの例 :モデルの評価.....	44

ストリームの例 :モデルの展開.....	44
----------------------	----

第 5 章 IBM Data Warehouse および IBM NetezzaAnalytics によるデータベース・

モデリング.....	45
スボス モデラー と、IBM Data Warehouse および IBM NetezzaAnalytics.....	45
統合の要件.....	46
統合の有効化.....	46
IBM NetezzaAnalytics または IBM Data Warehouse を構成する.....	46
IBM NetezzaAnalytics の ODBC ソースの作成.....	47
スボス モデラー で統合を有効にする.....	48
SQL の生成と最適化を有効にする.....	48
IBM NetezzaAnalytics および IBM Data Warehouse でのモデルの構築.....	48
フィールド・オプション.....	50
サーバー・オプション.....	50
モデル・オプション.....	50
モデルの管理.....	51
データベース・モデルの一覧表示.....	51
IBM Data WH 回帰ツリー.....	51
IBM Data WH 回帰ツリーの作成オプション - ツリーの成長.....	51
IBM Data WH ツリーの作成オプション - ツリーの剪定.....	52
Netezza 分裂クラスタリング.....	52
Netezza 分裂クラスタリングのフィールド・オプション.....	53
Netezza 分裂クラスタリングの作成オプション.....	53
IBM Data WH 一般化線型.....	54
IBM Data WH 一般化線形モデルのフィールド・オプション.....	54
IBM Data WH 一般化線形モデル・オプション - 全般.....	54
IBM Data WH 一般化線形モデル・オプション - 相互作用.....	55
IBM Data WH 一般化線形モデル・オプション - スコアリング・オプション.....	56
IBM Data WH ディシジョン・ツリー.....	56
インスタンスの重みとクラスの重み.....	56
Netezza ディシジョン・ツリーのフィールド・オプション.....	57
IBM Data WH ディシジョン・ツリーの作成オプション.....	57
IBM Data WH 線型.....	59
IBM Data WH 線型作成オプション.....	59
IBM Data WH KNN.....	59
IBM Data WH KNN モデル・オプション - 全般.....	59
IBM Data WH KNN モデル・オプション - スコアリング・オプション.....	60
IBM Data WH K-Means.....	60
IBM Data WH K-Means のフィールド・オプション.....	61
IBM Data WH K-Means の作成オプション・タブ.....	61
IBM Data WH Naive Bayes.....	62
Netezza ベイズ・ネットワーク.....	62
Netezza Bayes ネットワークのフィールド・オプション.....	62
Netezza Bayes ネットワークの作成オプション.....	62
Netezza 時系列.....	62
Netezza 時系列の値の補間.....	63
Netezza 時系列フィールド・オプション.....	65
Netezza 時系列構築オプション.....	65
Netezza 時系列のモデル・オプション.....	67
IBM Data WH TwoStep.....	67
IBM Data WH TwoStep フィールド・オプション.....	68
IBM Data WH TwoStep 作成オプション.....	68
IBM Data WH PCA.....	68
IBM Data WH PCA フィールド・オプション.....	69
IBM Data WH PCA 作成オプション.....	69
IBM Data WH モデルおよび Netezza モデルの管理.....	69
IBM Data Warehouse モデルと IBM NetezzaAnalytics モデルのスコアリング.....	69

IBM Data WH と Netezza のモデル・ナゲットの「サーバー」タブ.....	70
IBM Data WH ディシジョン・ツリー・モデル・ナゲット.....	70
IBM Data WH K-Means モデル・ナゲット.....	71
Netezza Bayes ネットワークのモデル・ナゲット.....	72
IBM Data WH Naive Bayes のモデル・ナゲット.....	73
IBM Data WH KNN モデル・ナゲット.....	73
Netezza 分裂クラスタリングのモデル・ナゲット.....	74
IBM Data WH PCA モデル・ナゲット.....	75
Netezza 回帰ツリー・モデル・ナゲット.....	75
IBM Data WH 線型モデル・ナゲット.....	76
Netezza 時系列モデル・ナゲット.....	76
IBM Data WH 一般化線型モデル・ナゲット.....	77
IBM Data WH TwoStep モデル・ナゲット.....	78
第 6 章 IBM Db2 for z/OS によるデータベース モデリング.....	79
IBM SPSS Modeler および IBM Db2 for z/OS.....	79
IBM Db2 for z/OS との統合の要件.....	79
IBM Db2 Analytics Accelerator for z/OS との統合の有効化.....	79
IBM Db2 for z/OS および IBM Analytics Accelerator for z/OS の構成.....	80
IBM Db2 for z/OS および IBM Db2 Analytics Accelerator の ODBC ソースの作成.....	80
IBM SPSS Modeler での IBM Db2 for z/OS の統合の有効化.....	80
SQL の生成と最適化を有効にする.....	81
IBM SPSS Modeler での IBM Db2 クライアントを使用した DSN の構成.....	81
IBM Db2 for z/OS でのモデル構築.....	82
IBM Db2 for z/OS モデル - フィールド オプション.....	83
IBM Db2 for z/OS モデル - サーバー・オプション.....	83
IBM Db2 for z/OS モデル - モデル・オプション.....	83
IBM Db2 for z/OS モデル - K-Means.....	83
IBM Db2 for z/OS モデル - K-Means のフィールド オプション.....	84
IBM Db2 for z/OS モデル - K-Means の作成オプション.....	84
IBM Db2 for z/OS モデル - Naive Bayes.....	84
IBM Db2 for z/OS モデル - ディシジョン・ツリー.....	85
IBM Db2 for z/OS モデル - ディシジョン ツリーのフィールド オプション.....	85
IBM Db2 for z/OS モデル - ディシジョン ツリーの作成オプション.....	85
IBM Db2 for z/OS モデル - ディシジョン ツリー ノード - クラスの重み.....	86
IBM Db2 for z/OS モデル - ディシジョン ツリー ノード - ツリーの剪定.....	86
IBM Db2 for z/OS モデル - 回帰ツリー.....	86
IBM Db2 for z/OS モデル - 回帰ツリーの作成オプション - ツリーの成長.....	87
IBM Db2 for z/OS モデル - 回帰ツリーの作成オプション - ツリーの剪定.....	87
IBM Db2 for z/OS モデル - TwoStep.....	88
IBM Db2 for z/OS モデル - TwoStep のフィールド オプション.....	88
IBM Db2 for z/OS モデル - TwoStep 作成オプション.....	88
IBM Db2 for z/OS モデル - TwoStep ナゲット - 「モデル」タブ.....	89
IBM Db2 for z/OS モデルの管理.....	89
IBM Db2 for z/OS モデルのスコアリング.....	89
IBM Db2 for z/OS ディシジョン・ツリー・モデル・ナゲット.....	89
IBM Db2 for z/OS K-Means モデル・ナゲット.....	90
IBM Db2 for z/OS Naive Bayes モデル ナゲット.....	90
IBM Db2 for z/OS 回帰ツリー モデル ナゲット.....	91
IBM Db2 for z/OS TwoStep モデル・ナゲット.....	91
特記事項.....	93
商標.....	94
製品資料に関するご使用条件.....	94
索引.....	97

前書き

IBM スポス モデラー は、IBM 企業強化用のデータ・マイニング・ワークベンチです。スポス モデラー を使用すると、企業はデータを詳しく調べることで、顧客および一般市民との関係を強化することができます。スポス モデラー から得られた情報を利用することで、大口の契約を見込める顧客を獲得し、クロスセルを提案できる案件を見つけだし、新たな顧客層を獲得し、不正を発見し、リスクを減少させ、政府機関へのサービスの提供を改善することができます。

スポス モデラー の視覚的インターフェースを使用すると、特定ビジネスの専門知識を適用し、より強力な予測モデルを実現し、解決までの時間を短縮します。スポス モデラー では、予測、分類、セグメンテーション、および関連性検出アルゴリズムなど、さまざまなモデル作成手法を提供しています。モデルを作成した後は、IBM スポスモデラー・ソリューション・パブリッシャにより、企業全体の意思決定者やデータベースにモデルを配布することが可能になります。

IBM Business Analytics について

IBM Business Analytics ソフトウェアは、意思決定者がビジネス・パフォーマンスの改善のために信頼できる完全で整合性があり、正確な情報を提供します。ビジネス・インテリジェンス、予測分析、財務実績および戦略管理、分析アプリケーションの包括的なポートフォリオを利用することによって、現在の実績を明確、迅速に理解し、将来の結果を予測することができます。豊富な業界ソリューション、効果が実証済みのプラクティス、熟練の専門サービスを組み合わせることにより、あらゆる規模の会社組織が、最高の生産性を達成し、信頼できる意思決定を自動化し、よりよい結果を実現させることができます。

このポートフォリオの一部として、IBM SPSS Predictive Analytics ソフトウェアを使用する組織は、将来のイベントを予測し、その洞察に基づいて積極的に行動し、より優れた業績を実現することができます。全世界の企業、政府、学術分野のお客様が IBM スポスの技術を活用し、不正行為を減少させ、リスクを軽減させながら、顧客の獲得、保持、成長において、競争優位を高めることができます。IBM スポス ソフトウェアを日常業務に組み込むことで、組織は企業としての予測能力を向上させ、意思決定の監督と自動化によってビジネス目標を達成し、高い競争力を獲得できます。詳細な情報、または営業担当者へのお問い合わせ方法については、<http://www.ibm.com/spss> を参照してください。

テクニカル・サポート

保守担当のお客様向けに技術サポートが提供されています。IBM 製品の使用方法、または対応するハードウェア環境へのインストールについてサポートが必要な場合は、テクニカル・サポートにご連絡ください。テクニカル・サポートのご利用には、<http://www.ibm.com/support> の IBM Web サイトをご覧ください。支援を要請される場合は、事前にユーザー、会社組織、およびサポート契約を明確にしておいていただくよう、お願いします。

第 1 章 IBM スポス モデラー について

IBM スポス モデラー は、ビジネスの専門知識を活用して予測モデルを迅速に作成したり、また作成したモデルをビジネス・オペレーションに展開して意思決定を改善できるようにする、一連のデータ・マイニング・ツールです。IBM スポス モデラー は業界標準の CRISP-DM モデルをベースに設計されたものであり、データ・マイニング・プロセス全体をサポートして、データに基づいてより良いビジネスの成果を達成できるようにします。

IBM スポス モデラー には、マシン学習、人工知能、および統計に基づいたさまざまなモデル作成方法が用意されています。「モデル作成」パレットを利用して、データから新しい情報を引き出したり、予測モデルを作成することができます。各手法によって、その長所や適した問題の種類が異なります。

スポス モデラー は、スタンドアロン製品として購入または スポスモデラー・サーバー と組み合わせてクライアントとして使用することができます。後のセクションで説明されているとおり、多くの追加オプションも使用することができます。詳しくは、<https://www.ibm.com/analytics/us/en/technology/spss/>を参照してください。

IBM スポス モデラー 製品

製品と関連するソフトウェアの IBM スポス モデラー ファミリーの構成は次のとおりです。

- IBM スポス モデラー
- IBM スポスモデラー・サーバー
- IBM スポス モデラー 管理コンソール (IBM スポス デプロイメント・マネージャー に付属)
- IBM スポス モデラー・バッチ
- IBM スポスモデラー・ソリューション・パブリッシャ
- IBM スポス コラボレーションおよびデプロイメント・サービス 用の IBM スポスモデラー・サーバー アダプター

IBM スポス モデラー

スポス モデラー はこの製品のすべての機能を搭載したバージョンであり、ユーザーのパーソナル・コンピューターにインストールし、そのコンピューターで実行します。スタンドアロン製品としてローカル・モードで スポス モデラー を実行するか、大規模なデータ・セットを使用する場合にパフォーマンスを向上させるために IBM スポスモデラー・サーバー と組み合わせて実行することができます。

スポス モデラー を使用して、プログラミングの必要なく、正確な予測モデルを迅速かつ直感的に構築することができます。独自のビジュアル・インターフェースを使用すると、データ・マイニング・プロセスを簡単に視覚化することができます。製品に組み込まれている高度な分析の支援を受けて、データ内に隠れたパターンやトレンドを発見することができます。結果をモデル化し、ビジネスチャンスを活用してリスクを軽減できるようになり、それらに影響を与える要因を理解することができます。

スポス モデラー は スポス モデラー・専門的 および スポスモデラー・保険料 の 2 つのエディションで使用できます。詳しくは、トピック 2 ページの『[IBM スポス モデラー のエディション](#)』を参照してください。

IBM スポスモデラー・サーバー

スポス モデラー は、クライアント/サーバー・アーキテクチャーを使用して、リソース集中型の操作が必要な要求を、強力なサーバー・ソフトウェアへ分散します。

スポスモデラー・サーバー は、1 つまたは複数の IBM スポス モデラー のインストールと組み合わせてサーバー・ホストで分散分析モードで継続的に実行する、別途ライセンスが必要な製品です。このように、スポスモデラー・サーバー では、メモリー集中型の操作を、クライアント コンピューターにデータをダウンロードせずにサーバー上で実行できるため、大きなデータ・セットで優れたパフォーマンスを発揮します。

IBM スポスマデラー・サーバー は、パフォーマンスと自動化のさらなる利点を提供し、SQL の最適化とデータベース内のモデリング機能をサポートしています。

IBM スポスマデラー 管理コンソール

モデラー管理コンソールは、スポスマデラー・サーバー 構成オプションの多くを管理するグラフィカル・ユーザー・インターフェースです。それらの構成オプションは、オプション・ファイルで設定することも可能です。コンソールは、IBM スポス デプロイメント・マネージャー に含まれています。コンソールを使用すると、スポスマデラー・サーバー インストール済み環境をモニターしたり、構成したりできます。スポスマデラー・サーバー の現在の顧客は、コンソールを無料で利用できます。アプリケーションは Windows コンピューターにのみインストールできますが、サポートされる任意のプラットフォームにインストールされたサーバーを管理できます。

IBM スポスマデラー・バッチ

データマイニングは、通常、対話型のプロセスですが、グラフィカル・ユーザー・インターフェースを必要とせずに、コマンドラインから スポスマデラー を実行することも可能です。例えば、ユーザーの介入なしで実行する長期実行または反復的なタスクがあります。スポスマデラー・バッチ は、通常のユーザー・インターフェースにアクセスせずに スポスマデラー の完全な分析機能のサポートを提供する製品の特別バージョンです。スポスマデラー・バッチ を使用するには、スポスマデラー・サーバー が必要です。

IBM スポスマデラー・ソリューション・パブリッシャー

スポスマデラー・ソリューション・パブリッシャーは、外部ランタイムで実行することができ、外部アプリケーションに埋め込まれる スポスマデラー ストリームのパッケージ版を作成することができるツールです。このように、スポスマデラー がインストールされていない環境で使用するための完全な スポスマデラー ストリームを公開して展開することができます。スポスマデラー・ソリューション・パブリッシャーは、個別のライセンスが必要とされている IBM スポス コラボレーションおよびデプロイメント・サービス - スコアリングサービスの一部として配布されています。このライセンスを使用すると、スポスマデラー・解決・パブリッシャー・ランタイム を受信し、公開されたストリームを実行することができます。

スポスマデラー・ソリューション・パブリッシャーについて詳しくは、IBM スポス コラボレーションおよびデプロイメント・サービスの資料を参照してください。IBM スポス コラボレーションおよびデプロイメント・サービスの IBM 資料には、『IBM SPSS Modeler Solution Publisher』および『IBM SPSS Analytics Toolkit』というセクションがあります。

IBM スポス コラボレーションおよびデプロイメント・サービス 用の IBM スポスマデラー・サーバー アダプター

さまざまな IBM スポス コラボレーションおよびデプロイメント・サービス 用のアダプターを使用すると、スポスマデラー および スポスマデラー・サーバー を IBM スポス コラボレーションおよびデプロイメント・サービス リポジトリとインタラクティブに機能させることができます。このように、リポジトリに展開された スポスマデラー ストリームは、複数のユーザーで共有したり、シンククライアントアプリケーションである IBM スポスマデラー・利点 からアクセスしたりできます。リポジトリをホストするシステムに、アダプターをインストールします。

IBM スポスマデラー のエディション

スポスマデラー は次のエディションで使用できます。

スポスマデラー・専門的

スポスマデラー・専門的 は、CRM システムで追跡する行動や対話、人口統計データ、購入行動や販売データなど、多くの構造化データを処理するために必要なすべてのツールを提供しています。

スボスマデラー・保険料

スボスマデラー・保険料 は、特化したデータ、または構造化されていないテキスト・データを処理するために スボスマデラー・専門的 を拡張する、別途ライセンスが必要な製品です。スボスマデラー・保険料 には、以下の IBMSPSS Modeler Text Analytics が含まれます。

IBMSPSS Modeler Text Analytics は、高度な言語技術と Natural Language Processing (NLP) を使用して、構造化されていない多様なテキスト・データをすばやく処理し、重要なコンセプトを抽出および組織化し、そしてそのコンセプトをカテゴリー別に分類します。抽出されたコンセプトとカテゴリーを、人口統計のような既存の構造化データと組み合わせ、IBM スボスマデラー の豊富なデータ・マイニング・ツールを適用する方法で、焦点を絞ったより良い決定を下すことができます。

IBM スボスマデラー・サブスクリプション

IBM スボスマデラー・サブスクリプション は、従来の IBM スボスマデラー クライアントとすべて同じ予測分析機能を提供します。Subscription エディションの場合、定期的に製品アップデートをダウンロードできます。

資料

資料は、スボスマデラーの **ヘルプ** メニューから入手できます。ここからオンラインの IBM 資料が開きます。この資料は、製品外部からいつでも利用できます。

各製品の完全な資料 (インストール手順を含む) は、<https://www.ibm.com/support/pages/spss-modeler-184-documentation> でも PDF 形式で入手できます。

スボスマデラー・専門的 ドキュメント

スボスマデラー・専門的 のドキュメントスイート (インストール手順を除く) は次のとおりです。

- **IBM スボスマデラー ユーザーズ・ガイド:** スボスマデラー の使用への全体的な入門で、データストリームの作成方法、欠損値の処理方法、CLEM 式の作成方法、プロジェクトおよびレポートの処理方法、および IBM スボスマデラー コラボレーションおよびデプロイメント・サービス または IBM スボスマデラー・利点に展開するためのストリームのパッケージ方法が含まれています。
- **IBM スボスマデラー 入力ノード、プロセス・ノード、出力ノード:** 各種形式のデータの読み取り、処理、および出力に使用するすべてのノードの説明です。これは、モデル作成ノード以外のすべてのノードについての説明です。
- **IBM スボスマデラー モデル作成ノード:** データ・マイニング・モデルの作成に使用するすべてのノードについての説明です。IBM スボスマデラー には、マシン学習、人工知能、および統計に基づいたさまざまなモデル作成方法が用意されています。
- **IBM スボスマデラー アプリケーション・ガイド:** このガイドの例では、特定のモデル作成手法および技法について、簡単に対象を絞って紹介します。本ガイドのオンラインバージョンは、「ヘルプ」メニューからも利用できます。詳しくは、トピック『』を参照してください。[4 ページの『アプリケーションの例』](#)
- **IBM スボスマデラー Python スクリプトとオートメーション:** Python スクリプトによるシステムの自動化に関する情報です。ノードおよびストリームの操作に使用できるプロパティを含めて説明します。
- **IBM スボスマデラー 展開ガイド:** IBM スボスマデラー デプロイメント・マネージャー のもとで処理されるジョブ内のステップとして IBM スボスマデラー のストリームを実行することに関する情報。
- **IBM スボスマデラー データベース内 マイニング・ガイド:** サード・パーティー製アルゴリズムを使用してご使用のデータベースの能力を利用してパフォーマンスを向上させ、分析機能の範囲を拡張する方法に関する情報を示します。
- **IBM スボスマデラー・サーバー 管理およびパフォーマンス・ガイド:** IBM スボスマデラー・サーバー の構成方法と管理方法に関する情報。
- **IBM スボスマデラー デプロイメント・マネージャー ユーザー・ガイド:** IBM スボスマデラー・サーバー の監視や構成を行うための デプロイメント・マネージャー アプリケーションに組み込まれている管理コンソール・ユーザー・インターフェースの使用法に関する情報。

- **IBM スポス モデラー CRISP-DM ガイド:** スポス モデラー でのデータ・マイニングに対する CRISP-DM 方法の使用に関するステップバイステップのガイドです。
- **IBM スポス モデラー・バッチ ユーザーズ・ガイド:** IBM スポス モデラー をバッチ・モードで使用するための完全ガイドで、バッチ・モードでの実行およびコマンド・ライン引数の詳細について説明します。このガイドは、PDF 形式のみです。

スポスモデラー・保険料 ドキュメント

スポスモデラー・保険料 のドキュメント スイート (インストール手順を除く) は次のとおりです。

- **SPSSModeler Text Analytics ユーザーズ・ガイド:** スポス モデラー でテキスト分析を使用する場合の情報。テキスト・マイニング・ノード、インタラクティブ・ワークベンチ、テンプレートなどについて説明します。

アプリケーションの例

スポス モデラー のデータ・マイニング・ツールは、多様なビジネスおよび組織の問題解決を支援しますが、アプリケーションの例では、特定のモデル作成手法および技術に関する簡単で、目的に沿った説明を行います。ここで使用するデータ・セットは、データ・マイニング担当者が管理するような大規模データ・ストアと比較すると非常に小規模ですが、関係する概念および手法は実際のアプリケーションにも拡張できます。

その例を参照するには、スポス モデラー の「ヘルプ」メニューから「**アプリケーションの例**」をクリックしてください。

データ・ファイルとサンプル・ストリームは、製品のインストール・ディレクトリーの **Demos** フォルダーにインストールされています。詳しくは、4 ページの『**Demos フォルダー**』を参照してください。

データベース・モデル作成の例: 例は、『**IBM スポス モデラー データベース内マイニング・ガイド**』を参照してください。

スクリプトの例: 例は、『**IBM スポス モデラー スクリプトとオートメーションガイド**』を参照してください。

Demos フォルダー

アプリケーション・サンプルで使用されるデータ・ファイルとサンプル・ストリームは、製品インストール・ディレクトリー (例: C:\Program Files\IBM\SPSS\Modeler\<version>\Demos) の下の **Demos** フォルダーにインストールされます。このフォルダーには、Windows の「スタート」メニューの **IBM スポス モデラー プログラム・グループ** から、または **ファイル > ストリームを開く** ダイアログ・ボックスの最近使用したディレクトリーのリストで **Demos** をクリックしてアクセスすることもできます。

ライセンスの追跡

スポス モデラー を使用すると、ライセンスの使用状況が一定の間隔で追跡され、ログに記録されます。ログに記録されるライセンス メトリックは **AUTHORIZED_USER** と **CONCURRENT_USER** であり、ログに記録されるメトリックのタイプは、スポス モデラー に使用するライセンスのタイプによって決まります。

作成されたログ ファイルは IBM License Metric Tool によって処理可能であり、そのファイルからライセンス使用状況レポートを生成できます。

ライセンスログファイルは、スポス モデラー クライアントログファイルが記録されるディレクトリーと同じディレクトリーに作成されます (デフォルトでは %ALLUSERSPROFILE%\IBM\SPSS\Modeler/<version>/log)。

第2章 データベース内マイニング

データベース・モデル作成の概要

IBM スポスモデラー・サーバー は、データベース・ベンダーから入手できる、IBM Netezza、Oracle Data Miner、Microsoft Analysis Services などのデータ・マイニングおよびモデル作成のツールとの統合をサポートしています。データベース内でのモデルの作成、スコアリング、および格納は、すべて IBM スポス モデラー アプリケーション内から実行できます。これにより、ベンダーが提供するデータベース固有のアルゴリズムを活用しながら、IBM スポス モデラー の分析機能と使いやすさを、データベースのパワーとパフォーマンスに組み合わせることが可能となります。モデルはデータベース内部に構築され、通常の方法で IBM スポス モデラー インターフェースによって参照、スコアリングすることができます。さらに必要に応じて IBM スポスモデラー・ソリューション・パブリッシャを使用して展開することもできます。使用可能なアルゴリズムは、IBM スポス モデラー の「データベースモデル作成」パレットにあります。

データベースのネイティブなアルゴリズムへのアクセスに IBM スポス モデラー を使用することにはいくつかの利点があります。

- データベース内アルゴリズムはデータベース・サーバーと密接に統合されることが多く、性能が向上することがあります。
- 「データベース内」で作成され格納されるモデルは、展開やデータベースにアクセス可能なアプリケーションとの共有をより簡単に実行できる場合があります。

SQL 生成。 データベース内モデル作成は SQL 生成と異なり、「SQL プッシュバック」として知られています。この機能を使用すると、パフォーマンス向上のためにデータベースに「プッシュバック」できる（つまり実行する）ネイティブの IBM スポス モデラー 操作の SQL 分を生成することができます。例えば、レコード結合、レコード集計、条件抽出のノードは、すべて、この方法でデータベースにプッシュバックできる SQL コードを生成します。SQL 生成は、データベース・モデル作成と組み合わせて使用すると、データベースの開始から終了までの実行が可能なストリームを生み出す場合があります。同時に、IBM スポス モデラー で実行するストリームに大幅なパフォーマンスの向上をもたらします。

注： DB モデリングおよび SQL 最適化では、IBM スポスモデラー・サーバー 接続が IBM スポス モデラー コンピュータ上で有効になっている必要があります。この設定を有効にすると、データベース・アルゴリズムにアクセスし、IBM スポス モデラー から SQL を直接プッシュバック、IBM スポスモデラー・サーバー にアクセスできます。現在のライセンスステータスを確認するには、IBM スポス モデラー メニューから次を選択します。

ヘルプ > 製品情報 > 追加の詳細

接続が有効な場合、「ライセンス ステータス」タブにオプション「サーバーの有効化」が表示されます。

サポートされているアルゴリズムの詳細は、この後の特定のベンダーの節を参照してください。

必要事項

DB モデリングを行う前に、以下の設定が必要です。

- 必要な分析コンポーネント (Microsoft Analysis Services または Oracle Data Miner) をインストールした適切なデータベースへの ODBC 接続。
- IBM スポス モデラー では、「ヘルパーアプリケーション」ダイアログ・ボックス（「ツール」>「ヘルパーアプリケーション」）でデータベース・モデリングを有効にする必要があります。
- IBM スポスモデラー・サーバー (使用している場合) と同様に、IBM スポス モデラー でも「ユーザー オプション」ダイアログ・ボックスで「SQL 生成」および「SQL 最適化」の設定を有効にする必要があります。SQL 最適化は、厳密には DB モデリングを使用するためには必要ではありませんが、パフォーマンス上の理由から使用することを強くお勧めします。

注： DB モデリングおよび SQL 最適化では、IBM スポスモデラー・サーバー 接続が IBM スポス モデラー コンピュータ上で有効になっている必要があります。この設定を有効にすると、データベース・アルゴリズムにアクセスし、IBM スポス モデラー から SQL を直接プッシュバック、IBM スポスモデラー・サーバー

にアクセスできます。現在のライセンスステータスを確認するには、IBM スポス モデラー メニューから次を選択します。

ヘルプ > 製品情報 > 追加の詳細

接続が有効な場合、「ライセンス ステータス」タブにオプション「**サーバーの有効化**」が表示されます。

詳細は、この後の特定ベンダーの節を参照してください。

モデルの構築

データベースのアルゴリズムを使用するモデルの構築とスコアリングのプロセスは、IBM スポス モデラーのほかの種類のデータマイニングに似ています。ノードとモデル作成「ナゲット」で作業の一般的なプロセスは、IBM スポス モデラー で作業する場合の他のストリームに似ています。唯一の相違点は、実際の処理とモデルの構築がデータベースに先送りされることです。

データベース・モデリング・ストリームは、概念的には IBM スポス モデラー 内の他のデータ・ストリームと同じです。ただし、このストリームは、Microsoft Decision Tree ノードを使用したモデルの構築などを含め、データベース内のすべての操作を実行します。このストリームを実行すると、IBM スポス モデラー は、データベースに構築と、作成されたモデルの保存を指示し、また詳細が IBM スポス モデラー にダウンロードされます。データベース内実行は、ストリーム内で紫色の影が付いているノードを使用することで示されます。

データの準備

データベース ネイティブ・アルゴリズムが使用されているか否かにかかわらず、データの準備をできるだけデータベースに先送りして、パフォーマンスを向上させる必要があります。

- その目的は、元のデータがデータベースに保存されている場合に、必要な上流操作がすべて SQL に変換できることを確認したうえで、データをデータベース内に保持しておくことです。これによって、データの IBM スポス モデラー へのダウンロードが防止されるため、ゲインを無効にする可能性があるボトルネックが回避され、ストリーム全体がデータベースで実行できるようになります。
- 元のデータが、データベースに保存されていない場合でも、DB モデリングを使用できます。この場合、データの準備は IBM スポス モデラー で行われ、準備が完了したデータ・セットは、モデル構築のために自動的にデータベースにアップロードされます。

モデル・スコアリング

データベース内マイニングを使用して IBM スポス モデラー から生成されるモデルは、通常の IBM スポス モデラー モデルとは異なります。それらは、生成されたモデル「ナゲット」として、モデル・マネージャーに表示されますが、実際には、リモートのデータ・マイニングやデータベース・サーバーに保持されるリモート・モデルです。IBM スポス モデラー に表示されるものは、単にこのリモート・モデルへの参照です。つまり、表示される IBM スポス モデラー モデルは、データベース・サーバー・ホスト名、データベース名およびモデル名などの情報を含む「hollow」モデルです。これは、データベース ネイティブ・アルゴリズムを使用して作成されるモデルを参照したりスコアリングしたりする場合に、理解すべき重要な違いです。

モデルが作成されると、IBM スポス モデラー で生成される他のモデルのようにスコアリングのため、ストリームに追加できます。アップストリーム操作が実行されていない場合でも、すべてのスコアリングがデータベース内で実行されます。(パフォーマンスを向上させるために、可能な場合はアップストリーム操作をデータベースにプッシュバックすることができますが、これはスコアリングを実行するための要件ではありません。)ほとんどの場合、データベース・ベンダーが提供する標準的なブラウザーを使用して、作成されたモデルをブラウズすることもできます。

参照する場合もスコアリングする場合も、Oracle Data Miner または Microsoft Analysis Services が動作しているサーバーへのライブ接続が必要です。

結果の表示と設定の指定

スコアリングの結果を表示し設定を指定するには、ストリーム領域のモデルをダブルクリックします。代わりに、モデルを右クリックして「参照」または「編集」を選択することもできます。特定のモデルのタイプに固有な設定

データベース・モデルのエクスポートと保存

「ファイル」メニューにあるオプションを使用すると、データベース・モデルは、IBM スポス モデラー で作成された他のモデルと同様の方法でモデル・ブラウザーからエクスポートできます。

1. モデル・ブラウザーの「ファイル」メニューから次のオプションを選択します。
 - 「テキストのエクスポート」を選択してモデル要約をテキスト・ファイルにエクスポート
 - 「HTML 形式でエクスポート」を選択してモデル要約を HTML ファイルにエクスポート
 - 「PMML をエクスポート」(IBM Db2 IM モデルのみをサポート)を選択してモデルを予測モデル・マークアップ言語 (PMML) としてエクスポート。このファイルは、他の PMML 互換ソフトウェアで使用できます。

注: 生成されたモデルは、「ファイル」メニューから「ノードの保存」を選択して保存することもできます。

モデルの整合性

生成された各データベース・モデルについて、IBM スポス モデラー は、そのデータベースに格納されているものと同じ名前のモデルへの参照と共に、モデル構造の説明を保存します。生成されたモデルの「サーバー」タブは、データベース上にある実際のモデルに一致するユニークなキーを表示します。

IBM スポス モデラー はこのランダムに生成されたキーを、モデルの整合性が維持されているかどうかを検査するために使用します。このキーは、モデルの構築時にモデルの説明に保存されます。展開ストリームを実行する前に、キーが一致することを確認しておくといでしょう。

1. データベースに保存されているモデルの説明と IBM スポス モデラー により保存されているランダム・キーとを比較することで整合性を検査するには、「検査」ボタンをクリックします。そのデータベース・モデルが見つからないか、キーが一致しない場合、エラーが報告されます。

生成された SQL の表示とエクスポート

生成された SQL コードは、実行前にプレビューすることができ、デバッグに役立ちます。

第3章 Microsoft Analysis Services によるデータベース・モデリング

IBM スポス モデラー と Microsoft Analysis Services

IBM スポス モデラー は、Microsoft SQL Server Analysis Services との統合をサポートします。この機能は IBM スポス モデラー のモデル作成ノードとして実装され、「DB モデル作成」パレットから使用できます。パレットが表示されていない場合は、「ヘルパー アプリケーション」ダイアログ・ボックスから、「Microsoft」タブ上にある「MS Analysis Services の統合」を有効にしてアクティブ化できます。詳しくは、トピック 11 ページの『[Analysis Services との統合を有効にする](#)』を参照してください。

IBM スポス モデラー は、次の Analysis Services アルゴリズムの統合をサポートします。

- ディジション・ツリー
- クラスタリング
- アソシエーション・ルール
- Naive Bayes
- 線型
- ニューラル・ネットワーク
- ロジスティック回帰
- 時系列
- シーケンス・クラスタリング

次の図は、クライアントから、IBM スポスモデラー・サーバー がデータベース内マイニングを管理するサーバーへのデータの流れを説明しています。モデル作成は、Analysis Services を使用して実行されます。作成されたモデルは、Analysis Services により保存されます。モデルへの参照が、IBM スポス モデラー ストリーム内に維持されます。次に、モデルは Analysis Services から、スコアリングのために Microsoft SQL Server、または、IBM スポス モデラー のいずれかにダウンロードされます。

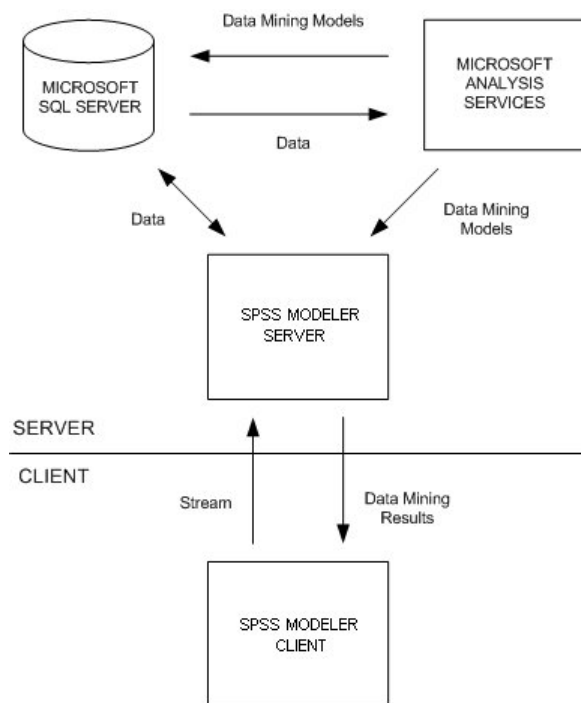


図 1. モデル作成中の IBM スポス モデラー、Microsoft SQL Server、および Microsoft Analysis Services 間のデータ・フロー

注：IBM スポスモデラー・サーバー は使用できますが、必須ではありません。IBM スポス モデラー クライアントはデータベース内マイニングの計算自体を処理することができます。

Microsoft Analysis Services との統合の要件

IBM スポス モデラー と共に Analysis Services アルゴリズムを使用してデータベース内モデル作成を行なう場合の前提条件を次に示します。場合によっては、これらの条件が満たされているかデータベース管理者に問い合わせて確認してください。

- Windows 上の IBM スポスモデラー・サーバー インストール済み環境 (分散モード) に対する IBM スポスモデラー の実行。Analysis Services との統合では、UNIX プラットフォームはサポートされません。

重要：IBM スポス モデラー ユーザーは、「IBM スポスモデラー・サーバー の追加の要件」の下に示す Microsoft の URL からダウンロードできる SQL Native Client ドライバーを使用して、ODBC 接続を構成する必要があります。IBM スポス データ・アクセス・パック に付属するドライバ (通常、IBM スポス モデラー で他の用途が推奨されている) は、この目的には推奨できません。ドライバは、「**統合 Windows 認証**」が有効になっている SQL Server を使用するように構成する必要があります。IBM スポス モデラー は、SQL Server 認証をサポートしていないためです。ODBC データ・ソースに関する質問がある場合は、データベース管理者に問い合わせてください。

- SQL Server がインストールされている必要があります。これは必ずしも IBM スポス モデラー と同一のホスト上の必要はありません。IBM スポス モデラー ユーザーには、データの読み込みおよび書き込み、テーブルとビューのドロップおよび作成に適切なアクセス権が必要です。

注：SQL Server Enterprise Edition が推奨されています。Enterprise Edition は、高度なパラメーターを提供することによって、アルゴリズムの結果を調整することができます。Standard Edition バージョンは、同一パラメーターを提供しますが、高度なパラメーターの中には編集できないものもあります。

- Microsoft SQL Server Analysis Services は、SQL Server と同一ホストにインストールされている必要があります。

IBM スポスモデラー・サーバー の追加の要件

IBM スポスマデラー・サーバー と共に Analysis Services のアルゴリズムを使用する場合は、次のコンポーネントが IBM スポスマデラー・サーバー ホスト・マシンにインストールされている必要があります。

注: SQL Server が IBM スポスマデラー・サーバー と同一ホストにインストールされている場合、これらのコンポーネントはすでに利用可能になっています。

- Microsoft SQL Server Analysis Services 10.0 OLE DB Provider (使用している OS に適切なバリエーションを選択してください)
- Microsoft SQL Server Native Client (使用している OS に適切なバリエーションを選択してください)
- Microsoft SQL Server 2008 または 2012 を使用している場合は Microsoft Core XML Services (MSXML) 6.0 も必要となる場合があります。

これらのコンポーネントをダウンロードするには、www.microsoft.com/downloads に移動して **.NET Framework** または (他のすべてのコンポーネントに) **SQL Server Feature Pack** を検索し、使用バージョンの SQL Server の最新パックを選択します。

これらは、ほかのパッケージが先にインストールされている必要がある場合があります。そのようなパッケージも、Microsoft Downloads Web サイトからダウンロードして利用できます。

IBM スポスマデラー の追加の要件

IBM スポスマデラー と共に Analysis Services アルゴリズムを使用するには、上記と同じコンポーネントおよびクライアントで以下をインストールする必要があります。

- Microsoft SQL Server Datamining Viewer Controls (使用している OS に適切なバリエーションを選択してください) も以下が必要です。
- Microsoft ADOMD.NET

これらのコンポーネントをダウンロードするには、www.microsoft.com/downloads に移動して **SQL Server Feature Pack** を検索し、使用バージョンの SQL Server の最新パックを選択します。

注: DB モデリングおよび SQL 最適化では、IBM スポスマデラー・サーバー 接続が IBM スポスマデラー コンピュータ上で有効になっている必要があります。この設定を有効にすると、データベース・アルゴリズムにアクセスし、IBM スポスマデラー から SQL を直接プッシュバック、IBM スポスマデラー・サーバー にアクセスできます。現在のライセンスステータスを確認するには、IBM スポスマデラー メニューから次を選択します。

ヘルプ > 製品情報 > 追加の詳細

接続が有効な場合、「ライセンス ステータス」タブにオプション「**サーバーの有効化**」が表示されます。

Analysis Services との統合を有効にする

Analysis Services との IBM スポスマデラー の統合を有効にするには、SQL Server および Analysis Services を設定し、ODBC ソースを作成して、IBM スポスマデラー の「ヘルパー アプリケーション」ダイアログ・ボックスで統合を有効にする必要があります。さらに、SQL の生成および最適化を有効にします。

注: Microsoft SQL Server および Microsoft Analysis Services が利用可能になっている必要があります。詳しくは、トピック [10 ページの『Microsoft Analysis Services との統合の要件』](#) を参照してください。

SQL Server を構成する

SQL Server を設定して、データベース内でスコアリングが実行できるようにします。

1. SQL Server ホスト・コンピューター上で次のレジストリ キーを作成します。

```
HKEY_LOCAL_MACHINE\SOFTWARE\Microsoft\MSSQLServer\Providers\MSOLAP
```

2. このキーに次の DWORD を追加します。

```
AllowInProcess 1
```

3. この変更が完了したら SQL Server を再起動します。

Analysis Services の設定

IBM スポス モデラー で Analysis Services との通信を開始する前に、まず、「Analysis Server Properties」ダイアログ・ボックスで 2 つの設定を手動で行う必要があります。

1. MS SQL Server Management Studio から Analysis Server へログインします。
 2. サーバー名を右クリックし「プロパティ」を選択して、「プロパティ」ダイアログ・ボックスにアクセスします。
 3. 「詳細プロパティを (すべて) 表示する」チェック・ボックスを選択します。
 4. 次のプロパティを変更します。
- DataMining\AllowAdHocOpenRowsetQueries の値を True に変更します (デフォルト値は False です)。
 - DataMining\AllowProvidersInOpenRowset の値を [all] に変更します (デフォルト値はありません)。

SQL Server の ODBC DSN を作成する

データベースを読み書きするには、ODBC データ・ソースがインストールされていて、該当するデータベースに対して必要に応じて読み取り権限や書き込み権限が設定されている必要があります。Microsoft SQL Native Client ODBC ドライバーが必要で、SQL Server で自動的にインストールされます。IBM スポス データ・アクセス・パック に付属するドライバ (通常、IBM スポス モデラー で他の用途が推奨されている) は、この目的には推奨できません。IBM スポス モデラー と SQL Server が異なるホストにある場合、Microsoft SQL Native Client ODBC ドライバーをダウンロードできます。詳しくは、トピック [10 ページの『Microsoft Analysis Services との統合の要件』](#) を参照してください。

ODBC データ・ソースに関する質問がある場合は、データベース管理者に問い合わせてください。

1. Microsoft SQL Native Client ODBC ドライバーを使用して、データ・マイニング・プロセスで使用される SQL Server データベースを指し示す ODBC DSN を作成します。その他は、デフォルトのドライバ設定を使用します。
 2. この DSN を使用するために、「**Integrated Windows Authentication と共に**」が選択されていることを確かめてください。
- IBM スポス モデラー と IBM スポスモデラー・サーバー が、それぞれ異なるコンピューター上で実行されている場合、両方のコンピューターに同じ ODBC DSN を作成します。それぞれのホストに同じ DSN 名が使われていることを確かめます。

IBM スポス モデラー で Analysis Services Integration を有効にする

IBM スポス モデラー を有効にして Analysis Services を使用するには、まず「ヘルパー アプリケーション」ダイアログ・ボックスでサーバー明細を準備する必要があります。

1. IBM スポス モデラー メニューから次の項目を選択します。
ツール > オプション > ヘルパー・アプリケーション
 2. 「Microsoft」タブをクリックします。
- **Microsoft Analysis Services 統合を有効化:** IBM スポス モデラー ウィンドウの下部にデータベース・モデリング・パレットを表示し、Analysis Services アルゴリズムのノードを追加します。
 - **Analysis サーバー・ホスト。** Analysis Services が実行されているマシンの名前を指定します。
 - **Analysis サーバー・データベース:** 「... (省略符号)」ボタンをクリックして使用可能なデータベースを選択できるサブダイアログボックスを開き、目的のデータベースを選択します。リストには、指定された Analysis サーバーで使用可能なデータベースが取り込まれます。Microsoft Analysis Services は名前付きデータベース内にデータ・マイニング・モデルを保管するため、IBM スポス モデラー によって作成された Microsoft モデルが保管される適切なデータベースを選択する必要があります。
 - **SQL Server Connection.** 分析サーバーに渡されるデータを保管するために SQL Server データベースによって使用される DSN 情報を指定します。Analysis Services データ・マイニング・モデルを構築するためのデータ提供に使用される、ODBC データ・ソースを選択します。フラット・ファイルまたは ODBC データ・ソースにより提供されるデータから Analysis Services モデルを構築する場合、この ODBC データ・ソースが指定している、データは SQL Server データベース内に作成される一時テーブルに自動的にアップロードされます。

- ・データ・マイニング・モデルを上書きしようとしているときに警告します。データベースに保管されているモデルが警告なしで IBM スポス モデラー によって上書きされないようにする場合に選択します。

注:「ヘルパー アプリケーション」ダイアログ・ボックスで行った設定は、さまざまな Analysis Services ノードの内部でオーバーライドできます。

SQL の生成と最適化を有効にする

1. IBM スポス モデラー メニューから次の項目を選択します。

ツール > ストリームのプロパティ > オプション

2. ナビゲーション ペインの「最適化」オプションをクリックします。
3. 「SQL 生成」オプションが有効になっていること確認します。この設定は、データベースのモデル作成が機能するために必要です。
4. 「SQL 生成の最適化」と「その他の実行を最適化」を選択します(絶対に必要な訳ではありませんが、最適化されたパフォーマンスを得るために、選択することを強くお勧めします)。

Analysis Services でのモデル構築

Analysis Services でモデルを構築するには、学習データ・セットを、SQL Server データベース内のテーブルまたはビューに配置する必要があります。データが、SQL Server がない場合、または、SQL Server では行えないデータ準備作業として IBM スポス モデラー で処理する必要がある場合、データはモデル構築の前に SQL Server の一時テーブルに自動的にアップロードされます。

Analysis Services モデルの管理

IBM スポス モデラー によって Analysis Services モデルを構築すると、IBM スポス モデラー 内にモデルが作成されると同時に、SQL Server データベース内でもモデルの作成や置き換えが行われます。この種類の IBM スポス モデラー モデルは、データベース・サーバーに格納されているデータベース・モデルの内容を参照します。IBM スポス モデラー は、IBM スポス モデラー モデルと SQL Server モデルの両方に、同一のモデル・キー文字列を生成して格納し、整合性チェックを実行します。



MS ディシジョン・ツリー モデル作成ノードは、カテゴリ属性と連続値属性の両方がある、予測的なモデル作成に使用されます。カテゴリ属性については、ノードで、データセット内の入力列間の関係に基づいた予測が行われます。例えば、どの顧客が自転車を購入するかを予測する計画で、若い顧客は 10 人のうち 9 人までが自転車を購入するのにに対し、年長者の顧客は 10 人のうち 2 人しか購入しない場合、ノードにより、年齢は自転車購入の予測に適した予測フィールドであると推論されます。ディシジョン・ツリーでは、このような傾向に基づいて特定の結果に向かう予測が行われます。連続値の属性については、ディシジョン・ツリーを分割する箇所を決めるために線型が使用されます。複数の列が予測値に設定される場合、または入力データが予測値に設定される入れ子のテーブルを含んでいる場合は、ノードにより、それぞれの列予測値に対応する個別のディシジョン・ツリーが作成されます。



MS クラスタリング モデル作成ノードでは、データセット内のケースをクラスター内で似たような特徴を持つグループにまとめる反復的な手法が使用されます。このようなグループ化は、データの検証、データ内の異常性の識別、予測の作成に役立ちます。クラスタリング・モデルで、通常の観察からは論理的に引き出せない可能性がある関係が、データセット内で識別されます。例えば、自転車で通勤する人は、通常は勤務先から遠い場所に住んでいないと論理的に認識できます。ただしアルゴリズムによって、自転車通勤者のそれほど明白とも言えない他の特徴も発見できます。クラスタリング・ノードは、対象フィールドが指定されないという点で、他のデータ・マイニング・ノードとは異なります。クラスタリング・ノードにより、データ内と、ノードに識別されるクラスターからの関係に従って、モデルが正確に調整されます。



Microsoft のアソシエーション・ルール モデル作成ノードは、購買アドバイス・エンジンにとって有用です。購買アドバイス・エンジンで、顧客がすでに購入したり興味を示したりしたアイテムに基づいて、商品が顧客に推奨されます。アソシエーション・ルールは、個々のケースとケースに含まれるアイテムの両方の識別子が含まれたデータセットに対して構築されます。ケース内のアイテムのグループは、**アイテムセット**といいます。アソシエーション・モデルは、一連のアイテムセットと、アイテムがケース内でグループ化される方法を記述したルールからなります。アルゴリズムで識別されるルールは、顧客のショッピングカート内にすでに存在するアイテムに基づいて、顧客の将来の購買を予測するのに使用できます。



MS Naive Bayes モデル作成ノードでは、対象フィールドと予測フィールド間の条件付きの確率が計算され、各列が独立していると見なされます。提示されたすべての予測変数は相互に依存関係がないものとして処理されるので、モデルは *naïve* と命名されます。この手法はほかの *Analysis Services* のアルゴリズムより計算上の強度が低く、そのため、モデル作成の準備段階で相互関係性を簡単に検出するのに役立ちます。このノードは、データを最初に検索し、その結果を適用して、計算に時間がかかってもより正確な結果が得られる他のノードと組み合わせて追加のモデルを作成するために使用できます。



MS 線型回帰 モデル作成ノードは、ディシジョン・ツリー・ノードのバリエーションです。ここで、**MINIMUM_LEAF_CASES** パラメーターは、ノードがマイニング・モデルの学習に使用するデータ・セット内のケースの総数以上に設定されます。パラメーターをこのように設定することによって、ノードは分岐を行うことはなく、線型を実行します。



MS Neural Network モデル作成ノードは、予測可能な属性の各状態が指定された場合、ニューラル・ネットワーク・ノードが入力属性の可能性のある各状態に対する可能性を計算するという点で、MS ディシジョン・ツリー・ノードと類似しています。後でこれらの可能性を使用し、入力属性に基づいて予測された属性の結果を予測することができます。



MS ロジスティック回帰 モデル作成ノードは、**HIDDEN_NODE_RATIO** パラメーターが 0 に設定されている MS ニューラル・ネットワーク・ノードのバリエーションです。この設定により、隠れ層を含まないニューラル・ネットワーク・モデルが作成されるため、ロジスティック回帰と同等になります。



MS Time Series モデル作成ノードでは、商品の売り上げなど、時間を経過した連続型値の予測に最適化された回帰アルゴリズムを提供します。ディシジョン・ツリーなど、その他の Microsoft アルゴリズムには、新しい情報の追加の列が傾向を予測するための入力として必要ですが、時系列モデルでは必要ありません。時系列モデルは、モデルの作成に使用される元のデータセットにのみ基づいて傾向を予測できます。予測を行う場合、モデルに新しいデータを追加して、自動的に傾向分析の新しいデータを結合します。詳しくは、トピック [16 ページの『MS Time Series ノード』](#)を参照してください。



MS Sequence Clustering モデル作成ノードは、データ内の順序が指定されたシーケンスを特定し、この分析の結果をクラスタリング手法によって結合し、シーケンスとその他の属性に基づいてクラスターを生成します。詳しくは、トピック [18 ページの『MS Sequence Clustering ノード』](#)を参照してください。

それぞれのノードは、IBM スポス モデラー ウィンドウの下部にある「DB モデリング」パレットからアクセスできます。

すべてのアルゴリズム・ノードに共通の設定

次の設定は、すべての Analysis Services アルゴリズムに共通です。

サーバー・オプション

「サーバー」タブでは Analysis サーバー・ホストとデータベース、SQL Server データ・ソースを設定します。ここで指定したオプションで、「ヘルパー アプリケーション」ダイアログソックスの「Microsoft」タブで指定した設定が上書きされます。詳しくは、トピック [11 ページの『Analysis Services との統合を有効にする』](#)を参照してください。

注：このタブの一種も、また、Analysis Services モデルをスコアリングする際に使用することができます。詳しくは、トピック [19 ページの『Analysis Services モデル・ナゲットの「サーバー」タブ』](#)を参照してください。

モデル オプション

最も基本的なモデルを構築するためには、作業の前に「モデル」タブの設定を指定する必要があります。「エキスパート」タブではスコアリング方法やその他の詳細設定を指定できます。

次の基本的なモデル作成オプションを利用できます。

モデル名。 ノードの実行時に作成されたモデルに割り当てられる名前を指定します。

- **自動** モデル名は、対象フィールドまたは ID フィールド名、または、クラスタリング・モデル・ノードのように対象が指定されていない場合モデル タイプの名前に基づいて自動的に生成されます。
- **カスタム。** 作成されたモデルのカスタム名を指定できるようになります。

データ区分データを使用。 現在のデータ区分フィールドに基づいて、データを、学習用、テスト用、および検証用の独立したサブセット、つまりサンプルに分割します。1 組のサンプルをモデルの作成に使用し、別の組のサンプルで生成したモデルをテストすることにより、そのモデルが、このデータに似た性質を持つより大きなデータセットにどの程度適用できるかについての目安を得ることができます。データ区分フィールドがストリーム中で指定されていない場合、このオプションは無視されます。

ドリルスルーを使用します。 このオプションが表示されている場合は、モデルを照会して、モデルに含まれているケースの詳細を確認できます。

固有フィールド。 ドロップダウン・リストから、各ケースを一意的に識別するフィールドを選択します。通常、これは「**CustomerID**」などの ID フィールドです。

MS ディジション・ツリーのエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

MS クラスタリングのエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

MS Naive Bayes のエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

MS Linear Regression のエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

MS Neural Network のエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

MS Logistic Regression エクスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

MS アソシエーション・ルール・ノード

MS のアソシエーション・ルール・モデル作成ノードは、購買アドバイス・エンジンにとって有用です。購買アドバイス・エンジンで、顧客がすでに購入したり興味を示したりしたアイテムに基づいて、商品が顧客に推奨されます。アソシエーション・ルールは、個々のケースとケースに含まれるアイテムの両方の識別子が含まれたデータセットに対して構築されます。ケース内のアイテムのグループは、**アイテムセット**といいます。

アソシエーション・モデルは、一連のアイテムセットと、アイテムがケース内でグループ化される方法を記述したルールからなります。アルゴリズムで識別されるルールは、顧客のショッピングカート内にすでに存在するアイテムに基づいて、顧客の将来の購買を予測するのに使用できます。

テーブル形式データの場合、アルゴリズムは生成された各推奨事項 (\$M-field) の確率 (\$MP-field) を示すスコアを作成します。トランザクション形式データの場合、アルゴリズムは生成された各推奨事項 (\$M-field) のサポート (\$MS-field)、確率 (\$MP-field) および調整済み確率 (\$MAP-field) のスコアを作成します。

要求

トランザクション形式アソシエーション・モデルの要件は次のとおりです。

- **固有フィールド。** 関連ルール・モデルには、レコードを一意的に識別するキーが必要です。
- **ID フィールド:** MS アソシエーション・ルール・モデルをトランザクション形式データで作成する場合、各トランザクションを特定する ID フィールドが必要です。ID フィールドは一意のフィールドと同じように設定できます。
- **少なくとも 1 つの入力フィールド。** 関連ルール・アルゴリズムには、少なくとも 1 つの入力フィールドが必要です。
- **対象フィールド:** MS アソシエーション・モデルをトランザクション形式データで作成する場合、対象フィールドはユーザーが購入した製品など、トランザクション・フィールドでなければなりません。

MS アソシエーション・ルールのエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

MS Time Series ノード

MS Time Series モデル作成ノードは、次の 2 つの予測をサポートしています。

- 将来
- 過去

将来の予測は、過去のデータ以降の指定された期間、対象フィールド値を推定し、これは常に実行されます。過去の予測は、過去のデータに実際の値が指定されている期間、対象フィールド値を推定します。**過去の予測**とは、過去のデータに実際の値がある指定した数の期間について、推定された対象フィールド値です。過去の予測を使用して、実際の過去の値と予測値を比較することで、モデルの品質を評価できます。予測の始点の値によって、過去の予測が実行されるかどうかが決まります。

IBM スポス モデラー 時系列ノードとは異なり、MS Time Series ノードの前に時間区分ノードは必要ありません。さらに異なるのは、デフォルトで時系列データのすべての過去の行ではなく、予測行にのみスコアが作成されるという点です。

要求

MS Time Series モデルの要件は次のとおりです。

- **単一の鍵時刻フィールド。** 各モデルには、モデルが使用するタイム・スライスを定義する、ケース・シリーズとして使用される 1 つの数値フィールドまたは日付フィールドが含まれている必要があります。キー時間フィールドのデータ型は、日付/時刻型または数値型です。ただし、フィールドには連続型の値が含まれている必要があります、各系列の値は一意でなければなりません。
- **単一の対象フィールド。** 各モデルに指定できる対象フィールドは 1 つだけです。対象フィールドのデータ型は連続型の値でなければなりません。例えば、収入、売上、温度などの数値属性が、時間を経てどのように変化するかを予測できます。ただし、購入状況または学歴のレベルなど、カテゴリー型値を含むフィールドを対象フィールドとして使用できません。
- **少なくとも 1 つの入力フィールド。** MS Time Series アルゴリズムには、少なくとも 1 つの入力フィールドが必要です。入力フィールドのデータ型は連続型の値でなければなりません。連続型以外の入力フィールドは、モデルの作成時に無視されます。
- **データ・セットはソートする必要があります。** ソートしない場合、エラーが発生してモデル作成が中断します。

MS Time Series のモデル・オプション

モデル名。 ノードの実行時に作成されたモデルに割り当てられる名前を指定します。

- **自動** モデル名は、対象フィールドまたは ID フィールド名、または、クラスタリング・モデル・ノードなどのように対象が指定されていない場合モデル タイプの名前に基づいて自動的に生成されます。
- **カスタム。** 作成されたモデルのカスタム名を指定できるようになります。

データ区分データを使用。 データ区分フィールドが定義されている場合、このオプションでは学習用データ区分からのデータのみがモデル構築に使用されるようにします。

ドリルスルーを使用します。 このオプションが表示されている場合は、モデルを照会して、モデルに含まれているケースの詳細を確認できます。

固有フィールド。 ドロップダウン・リストから、時系列モデルの作成に使用するキー時間フィールドを選択します。

MS Time Series のエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

過去の予測を行っている場合、スコアリング結果に含めることができる過去のステップ数は、 $(\text{HISTORIC_MODEL_COUNT} * \text{HISTORIC_MODEL_GAP})$ の値によって決まります。デフォルトでは、この制約は 10 で、過去の予測は 10 件だけ行われます。この場合、モデル・ナゲットの「設定」タブの「過去の予測」に -10 より小さい値を入力すると、エラーが発生します (20 ページの『MS Time Series モデル・ナゲットの「設定」タブ』を参照)。過去の予測についてもっと知りたい場合、HISTORIC_MODEL_COUNT または HISTORIC_MODEL_GAP の値を大きくすることができますが、モデルの作成時間が長くなります。

MS Time Series の設定オプション

推定の開始。 予測を開始する期間を指定します。

- **開始点: 新しい予測。** 将来の予測を開始する期間。過去のデータの最後の期間からのオフセットとして表されます。例えば、過去のデータが 12/99 に終了し、予測を 01/00 に開始したい場合、1 の値を使用します。ただし、予測を 03/00 に開始したい場合、3 の値を使用します。
- **開始: 過去の予測。** 過去の予測を開始する期間。過去のデータの最後の期間からの負のオフセットとして表されます。例えば、過去のデータが 12/99 で終了し、データの終わりの 5 つの期間の過去の予測を行う場合、-5 の値を使用します。

推定の終了。 予測を停止する期間を指定します。

- **予測の終了ステップ。** 予測を停止する期間。過去のデータの最後の期間からのオフセットとして表されます。例えば、過去のデータが 12/99 で終了し、予測を 6/00 で終了する場合、ここで 6 の値を使用します。将来の予測の場合、値は「始点」の値以上でなければなりません。

MS Sequence Clustering ノード

MS Sequence Clustering ノードでは、シーケンス分析アルゴリズムを使用して、次のパスまたはシーケンスでリンクできるイベントを含むデータを探索します。例えば、ユーザーが Web サイトをナビゲートまたは参照する場合に作成されるクリックパス、または顧客がオンライン ショップのショッピング カートに品物を追加する順序などとなります。アルゴリズムは、同一のシーケンスをグループ化またはクラスターリングして、最も一般的なシーケンスを検出します。

要求

Microsoft Sequence Clustering モデルの要件は次のとおりです。

- **ID フィールド:** Microsoft Sequence Clustering アルゴリズムには、シーケンス情報をトランザクション形式で保存する必要があります。このため、各トランザクションを特定する ID フィールドが必要になります。
- **少なくとも 1 つの入力フィールド。** アルゴリズムには少なくとも 1 つの入力フィールドが必要です。
- **シーケンス・フィールド。** 尺度は連続型でなければなりません。例えば、フィールドがシーケンス内のイベントを特定する限り、Web ページの ID、整数、または文字列を使用できます。各シーケンスにはシーケンス ID を 1 つだけ使用でき、また各モデルにシーケンスを 1 種類だけ使用できます。シーケンス・フィールドは、ID フィールドおよび一意フィールドとは別のフィールドでなければなりません。
- **対象フィールド:** 対象フィールドは、シーケンス・クラスターリング・モデルを作成する場合に必要です。
- **固有フィールド。** シーケンス・クラスターリング・モデルには、レコードを一意的に識別するキー・フィールドが必要です。一意のフィールドは ID フィールドと同じように設定できます。

MS Sequence Clustering のフィールド・オプション

すべてのモデル作成ノードには、「フィールド」タブがあり、そこからモデルの作成に使用するフィールドを指定できます。

シーケンス・クラスターリング・モデルを作成する前に、対象フィールドや入力フィールドを指定する必要があります。MS Sequence Clustering ノードの場合、上流のデータ型ノードのフィールド情報を使用できません。ここでフィールド設定を指定する必要があります。

ID。 リストから ID フィールドを選択します。ID フィールドとして使用できるのは、数値またはシンボル値のフィールドです。選択したフィールドでは、一意の値がそれぞれ、ある分析ユニットを示している必要があります。例えば、マーケット・バスケット分析なら、各 ID が 1 人の顧客を表します。Web ログ分析なら、各 ID が 1 台のコンピューター (IP アドレス) あるいは 1 人のユーザー (ログイン・データ) を表します。

入力: モデルの入力フィールドを選択してください。これらのフィールドには、シーケンス・モデル作成の対象となるイベントが含まれています。

順序。 シーケンス ID フィールドとして使用するフィールドをリストから選択します。例えば、フィールドがシーケンス内のイベントを特定する場合、Web ページの ID、整数、または文字列を使用できます。各シーケンスにはシーケンス ID を 1 つだけ使用でき、また各モデルにシーケンスを 1 種類だけ使用できます。シーケンス・フィールドは ID フィールド (このタブで指定) および一意フィールド (「モデル」タブで指定) とは異なるフィールドでなければなりません。

target. 対象フィールドとして使用するフィールド、値をシーケンス・データに基づいて予測しようとしているフィールドを選択します。

MS Sequence Clustering のエキスパート・オプション

「エキスパート」タブで利用可能なオプションは、選択したストリームの構造に応じて変化する可能性があります。各 Analysis Services モデル・ノードで選択したエキスパート・オプションの詳細は、ユーザー・インターフェースのフィールド レベルのヘルプを参照してください。

Analysis Services モデルのスコアリング

モデル・スコアリングが SQL Server 内に発生し Analysis Services で実行されます。データが IBM スポス モデラー 内にあったり、IBM スポス モデラー 内でデータを準備する必要がある場合は、データ・セットを一時テーブルにアップロードすることが必要な場合があります。データベース内マイニングを使用して IBM スポス モデラー から生成したモデルは、リモートのデータマイニング・サーバーまたはデータベース・サーバー上に保管されています。これは、Microsoft Analysis Services を使用して作成されたモデルを参照およびスコアリングする場合に、理解すべき重要な違いです。

IBM スポス モデラー では、通常、単一の予測および関連付けられている確率または確信度のみが提供されます。

モデル・スコアリング の例については、[21 ページの『Analysis Services マイニングの例』](#)を参照してください。

すべての Analysis Services モデルに共通の設定

次の設定は、すべての Analysis Services モデルに共通です。

Analysis Services モデル・ナゲットの「サーバー」タブ

「サーバー」タブで、データベース内マイニングに対する接続を指定します。このタブには、一意のモデル・キーも用意されています。このキーは、IBM スポス モデラー のモデル内と Analysis Services データベースに格納されているモデル・オブジェクトの説明内との両方で、モデルが作成され、格納されるときに、ランダムに生成されます。

「サーバー」タブでは Analysis サーバー・ホストとデータベース、スコアリング操作に対する SQL Server データ・ソースを設定します。ここで指定した設定は、IBM スポス モデラー の「ヘルパー アプリケーション」ダイアログ・ボックス または「モデル作成」ダイアログ・ボックスで指定した設定を上書きします。詳しくは、トピック [11 ページの『Analysis Services との統合を有効にする』](#)を参照してください。

モデル GUID。 モデル・キーを以下に示します。このキーは、IBM スポス モデラー のモデル内と Analysis Services データベースに格納されているモデル・オブジェクトの説明内との両方で、モデルが作成され、格納されるときに、ランダムに生成されます。

確認: このボタンをクリックして、Analysis Services データベースに保管されているモデルのキーに対してモデル・キーを検査します。これにより、モデルが Analysis サーバー内に存在し、モデルの構造に変化がないことを示すことを検証できます。

注: 「検査」ボタンはスコアリングに備えてストリーム領域に追加されたモデルにのみ使用できます。検査が失敗した場合、モデルが削除されたか、サーバー上の他のモデルにより置き換えられたかを調べます。

表示。 クリックすると、ディシジョン・ツリー・モデルがグラフィカルに表示されます。Decision Tree Viewer は IBM スポス モデラー のすべてのディシジョン・ツリー・アルゴリズムが共有し、その機能は同一です。

Analysis Services モデル・ナゲットの「要約」タブ

モデル・ナゲットの「要約」タブで、モデルそのもの(精度分析)、モデルで使用するフィールド(フィールド)、モデルの構築時に使用する設定(構築の設定)、およびモデルの学習(学習の要約)についての情報を表示します。

ノードを初めて参照する場合、「要約」タブの結果は閉じられています。目的の結果を表示するには、項目の左側にある展開コントロールを使用して項目を展開するか、または「**すべて展開**」ボタンをクリックしてすべての結果を表示します。見終わった結果を隠すには、展開コントロールを使用して目的の結果を省略するか、または「**すべて閉じる**」ボタンをクリックしてすべての結果を非表示にします。

分析。 特定のモデルについての情報を表示します。このモデル・ナゲットに接続されている精度分析ノードを実行した場合、その精度分析情報もこのセクションに表示されます。

フィールド。 対象フィールドおよびモデル構築時の入力として使われるフィールドが表示されます。

構築の設定: モデル構築時に使われる設定情報が表示されます。

学習の要約: モデルの種類、モデルの作成に使われたストリーム、モデルの作成者、モデルの作成日時、およびモデルの構築時間などの情報が表示されます。

MS Time Series モデル・ナゲット

MS Time Series モデルは、過去のデータではなく、予測された期間のみのスコアを生成します。

モデルに追加されるフィールドを次の表に示します。

表 1. モデルに追加されるフィールド	
フィールド名	説明
\$M-field	フィールドの予測値
\$Var-field	計算されたフィールドの分散
\$Stdev-field	フィールドの標準偏差

MS Time Series モデル・ナゲットの「サーバー」タブ

「サーバー」タブで、データベース内マイニングに対する接続を指定します。このタブには、一意のモデル・キーも用意されています。このキーは、IBM スポス モデラー のモデル内と Analysis Services データベースに格納されているモデル・オブジェクトの説明内との両方で、モデルが作成され、格納されるときに、ランダムに生成されます。

「サーバー」タブでは Analysis サーバー・ホストとデータベース、スコアリング操作に対する SQL Server データ・ソースを設定します。ここで指定した設定は、IBM スポス モデラー の「ヘルパー アプリケーション」ダイアログ・ボックスまたは「モデル作成」ダイアログ・ボックスで指定した設定を上書きします。詳しくは、トピック [11 ページの『Analysis Services との統合を有効にする』](#)を参照してください。

モデル GUID。 モデル・キーを以下に示します。このキーは、IBM スポス モデラー のモデル内と Analysis Services データベースに格納されているモデル・オブジェクトの説明内との両方で、モデルが作成され、格納されるときに、ランダムに生成されます。

確認: このボタンをクリックして、Analysis Services データベースに保管されているモデルのキーに対してモデル・キーを検査します。これにより、モデルが Analysis サーバー内に存在し、モデルの構造に変化がないことを示すことを検証できます。

注: 「検査」ボタンはスコアリングに備えてストリーム領域に追加されたモデルにのみ使用できます。検査が失敗した場合、モデルが削除されたか、サーバー上の他のモデルにより置き換えられたかを調べます。

表示。 クリックすると、時系列モデルがグラフィカルに表示されます。Analysis Services は、完成したモデルをツリーで表示します。対象フィールドの時を経た過去の値と、将来の予測値をともに表示するグラフとを表示することができます。

詳細は、<http://msdn.microsoft.com/en-us/library/ms175331.aspx> の MSDN ライブラリーのタイム シリーズ・ビューアーの説明を参照してください。

MS Time Series モデル・ナゲットの「設定」タブ

推定の開始。 予測を開始する期間を指定します。

- **開始点: 新しい予測。** 将来の予測を開始する期間。過去のデータの最後の期間からのオフセットとして表されます。例えば、過去のデータが 12/99 に終了し、予測を 01/00 に開始したい場合、1 の値を使用します。ただし、予測を 03/00 に開始したい場合、3 の値を使用します。
- **開始: 過去の予測。** 過去の予測を開始する期間。過去のデータの最後の期間からの負のオフセットとして表されます。例えば、過去のデータが 12/99 で終了し、データの終わりの 5 つの期間の過去の予測を行う場合、-5 の値を使用します。

推定の終了。 予測を停止する期間を指定します。

- **予測の終了ステップ。** 予測を停止する期間。過去のデータの最後の期間からのオフセットとして表されます。例えば、過去のデータが 12/99 で終了し、予測を 6/00 で終了する場合、ここで 6 の値を使用します。将来の予測の場合、値は「始点」の値以上でなければなりません。

MS Sequence Clustering のモデル・ナゲット

MS Sequence Clustering モデルに追加されるフィールドを次の表に示します (ここで、フィールド とは対象フィールドの名前です)。

表 2. モデルに追加されるフィールド	
フィールド名	説明
\$MC-field	シーケンスが属するクラスターの予測。
\$MCP-field	このシーケンスが予測クラスターに含まれる確率。
\$MS-field	フィールドの予測値
\$MSP-field	\$MS-field 値が適切である確率。

モデルのエクスポートとノードの生成

モデル集計をエクスポートして、テキストおよび HTML 形式のファイルに構造化できます。適切な場所に、適切な条件抽出ノード、フィルター・ノードを生成できます。

IBM スポス モデラー のモデル・ナゲットと同様、Microsoft Analysis Services モデル・ナゲットはレコードの直接生成とフィールド設定ノードをサポートします。モデル・ナゲットの「ノードの生成」メニューのオプションを使用して、次のノードを生成できます。

- 条件抽出ノード (項目が「モデル」タブで選択された場合のみ)
- フィルター・ノード

Analysis Services マイニングの例

いくつものサンプル・ストリームがあり、それらが IBM スポス モデラー を使用して MS Analysis Services データ・マイニングを行う方法を説明します。これらのストリームは、次の IBM スポス モデラー のインストール・フォルダーにあります。

¥Demos¥Database_Modelling¥Microsoft

注: Demos フォルダーには、Windows の「スタート」メニューの IBM スポス モデラー プログラム・グループからアクセスできます。

ストリームの例:ディシジョン・ツリー

次のストリームは、MS Analysis Services が提供するディシジョン・ツリー・アルゴリズムを使用するデータベース・マイニング・プロセスの例として、順番に使用されます。

表 3. ディジジョン・ツリー - ストリーム例	
ストリーム	説明
1_upload_data.str	フラット・ファイルのデータを整理して、データベースへアップロードするために使用されます。
2_explore_data.str	IBM スポス モデラー でのデータ探索の例として使用されます。
3_build_model.str	データベース固有のアルゴリズムを使用したモデルを構築します。
4_evaluate_model.str	IBM スポス モデラー でのモデル評価の例として使用されます。
5_deploy_model.str	データベース内スコアリングのためにモデルを展開します。

注: サンプルを実行するには、ストリームを順番に実行する必要があります。さらに、各ストリーム中の入力およびモデル作成ノードは、使用するデータベースの有効なデータ・ソースを参照するように更新する必要があります。

サンプル・ストリームで使用されるデータセットは、クレジットカード申請に関するものであり、カテゴリ型および連続型予測フィールドの混在について、分類上の問題を提示します。このデータ・セットの詳細は、サンプル・ストリームと同じフォルダーにある *crx.names* ファイルを参照してください。

このデータセットは、<ftp://ftp.ics.uci.edu/pub/machine-learning-databases/credit-screening/> にある UCI Machine Learning Repository から入手可能です。

ストリームの例: データのアップロード

最初のストリーム例、1_upload_data.str は、フラット・ファイルのデータを整理して SQL Server へアップロードするのに使用されます。

Analysis Services データ・マイニングは、キー・フィールドを必要とするため、この初期ストリームは、フィールド作成ノードを使用して、一意の値 1、2、3 を持つ *KEY* という名前の新しいフィールドをデータ・セットに追加します。これには、IBM スポス モデラー の @INDEX 関数が使用されます。

その次にある置換ノードは、欠損値の処理に使用され、*crx.data* テキスト・ファイルから読み込まれた空のフィールドを *NULL* 値で置き換えます。

ストリームの例: データの調査

2 番目の例のストリーム、2_explore_data.str を使用して、要約統計およびグラフなど、データの概要を取得するデータ検査ノードの使用方法を説明します。

データ検査レポート内のグラフをダブルクリックすると、指定されたフィールドについて、より深く検索した結果を表す詳細なグラフが生成されます。

ストリームの例: モデルの作成

3 番目のストリーム例、3_build_model.str では、IBM スポス モデラー でのモデル構築を説明します。ストリームにデータベース・モデルを追加したら、追加したモデルをダブルクリックして構築の設定を指定できます。

このダイアログ・ボックスの「モデル」タブでは、次の設定ができます。

1. 「キー」 フィールドを一意的 ID フィールドとして指定します。

「エキスパート」タブでは、モデル構築の設定を細かく指定できます。

実行する前に、モデル構築用に正しいデータベースが指定されていることを確認してください。「サーバー」タブを使用すると、設定を変更できます。

ストリームの例:モデルの評価

第4のストリーム例、`4_evaluate_model.str`では、IBM スポス モデラー を使用したデータベース内モデル作成の利点を説明します。モデルを実行すると、その結果をユーザーのデータストリームに追加したり、IBM スポス モデラー が提供するいくつかのツールを使用してモデルを評価できます。

モデル作成結果の表示

モデル・ナゲットをダブルクリックすると、結果を確認できます。「要約」タブには、結果を表示するルールツリー・ビューがあります。また、「サーバー」タブにある「**ビュー**」ボタンをクリックして、ディシジョン・ツリー・モデルのグラフ表示を行うことができます。

モデル作成結果の評価

サンプル・ストリームの精度分析ノードを使用すると、予測フィールドとその対象フィールド間の一致パターンを表す一致行列を作成できます。結果を表示するには精度分析ノードを実行します。

サンプル・ストリームの評価ノードを使用して、このモデルによりどの程度、精度が改善されたかを示すゲイン・グラフを作成できます。結果を表示するには評価ノードを実行します。

ストリームの例:モデルの展開

モデルの精度が満足できるものになったら、外部のアプリケーションと共に使用するために展開したり、データベースに保存するために発行できます。最後のストリームの例、`5_deploy_model.str`では、データは表 CREDIT から読み込まれ、スコアリングされてから、データベース・エクスポート・ノードを使用して、表 CREDITSCORES に発行されます。

ストリームを実行すると次の SQL が生成されます。

```
DROP TABLE CREDITSCORES

CREATE TABLE CREDITSCORES ( "field1" varchar(1),"field2" varchar(255),"field3" float,"field4" varchar(1),"field5" varchar(2),"field6" varchar(2),"field7" varchar(2),"field8" float,"field9" varchar(1),"field10" varchar(1),"field11" int,"field12" varchar(1),"field13" varchar(1),"field14" int,"field15" int,"field16" varchar(1),"KEY" int,"$M-field16" varchar(9),"$MC-field16" float )

INSERT INTO CREDITSCORES ("field1","field2","field3","field4","field5","field6","field7","field8","field9","field10","field11","field12","field13","field14","field15","field16","KEY","$M-field16","$MC-field16")
SELECT T0.C0 AS C0,T0.C1 AS C1,T0.C2 AS C2,T0.C3 AS C3,T0.C4 AS C4,T0.C5 AS C5,
T0.C6 AS C6,T0.C7 AS C7,T0.C8 AS C8,T0.C9 AS C9,T0.C10 AS C10,
T0.C11 AS C11,T0.C12 AS C12,T0.C13 AS C13,T0.C14 AS C14,
T0.C15 AS C15,T0.C16 AS C16,T0.C17 AS C17,T0.C18 AS C18
FROM (
SELECT CONVERT(NVARCHAR,[TA].[field1]) AS C0, CONVERT(NVARCHAR,[TA].[field2]) AS C1,
[TA].[field3] AS C2, CONVERT(NVARCHAR,[TA].[field4]) AS C3,
CONVERT(NVARCHAR,[TA].[field5]) AS C4, CONVERT(NVARCHAR,[TA].[field6]) AS C5,
CONVERT(NVARCHAR,[TA].[field7]) AS C6, [TA].[field8] AS C7,
CONVERT(NVARCHAR,[TA].[field9]) AS C8, CONVERT(NVARCHAR,[TA].[field10]) AS C9,
[TA].[field11] AS C10, CONVERT(NVARCHAR,[TA].[field12]) AS C11,
CONVERT(NVARCHAR,[TA].[field13]) AS C12, [TA].[field14] AS C13,
[TA].[field15] AS C14, CONVERT(NVARCHAR,[TA].[field16]) AS C15,
[TA].[KEY] AS C16, CONVERT(NVARCHAR,[TA].[$M-field16]) AS C17,
[TA].[$MC-field16] AS C18
FROM openrowset('MSOLAP',
'Datasource=localhost;Initial catalog=FoodMart 2000',
'SELECT [T].[C0] AS [field1],[T].[C1] AS [field2],[T].[C2] AS [field3],
[T].[C3] AS [field4],[T].[C4] AS [field5],[T].[C5] AS [field6],
[T].[C6] AS [field7],[T].[C7] AS [field8],[T].[C8] AS [field9],
[T].[C9] AS [field10],[T].[C10] AS [field11],[T].[C11] AS [field12],
[T].[C12] AS [field13],[T].[C13] AS [field14],[T].[C14] AS [field15],
[T].[C15] AS [field16],[T].[C16] AS [KEY],[CREDIT1].[field16] AS [$M-field16],
PredictProbability([CREDIT1].[field16]) AS [$MC-field16]
FROM [CREDIT1] PREDICTION JOIN
openrowset('MSDASQL',
'Dsn=LocalServer;Uid=pwd=','SELECT T0."field1" AS C0,T0."field2" AS C1,
T0."field3" AS C2,T0."field4" AS C3,T0."field5" AS C4,T0."field6" AS C5,
T0."field7" AS C6,T0."field8" AS C7,T0."field9" AS C8,T0."field10" AS C9,
T0."field11" AS C10,T0."field12" AS C11,T0."field13" AS C12,
T0."field14" AS C13,T0."field15" AS C14,T0."field16" AS C15,
T0."KEY" AS C16 FROM "dbo".CREDITDATA T0')) AS [T]
ON [T].[C2] = [CREDIT1].[field3] and [T].[C7] = [CREDIT1].[field8]
and [T].[C8] = [CREDIT1].[field9] and [T].[C9] = [CREDIT1].[field10]
and [T].[C10] = [CREDIT1].[field11] and [T].[C11] = [CREDIT1].[field12]
and [T].[C14] = [CREDIT1].[field15]') AS [TA]
) T0
```

第 4 章 Oracle Data Mining によるデータベース・モデリング

Oracle Data Mining について

IBM スポス モデラー は、Oracle Data Mining (ODM) との統合をサポートします。このリリースでは、Oracle RDBMS 内にデータ・マイニング・アルゴリズム・ファミリーが密接に組み込まれています。これらの機能は、IBM スポス モデラー のグラフィカル・ユーザー・インターフェースとワークフロー指向の開発環境でアクセスでき、顧客は ODM が提供するデータ・マイニング・アルゴリズムを使用できます。

IBM スポス モデラー は、Oracle Data Mining の次のアルゴリズムを統合できます。

- Naive Bayes
- Adaptive Bayes
- Support Vector Machine (SVM)
- 一般化線型モデル (GLM)*
- ディジション・ツリー
- O-Cluster
- K-means
- 非負数マトリックス因数分解 (NMF)
- Apriori
- Minimum Descriptor Length (MDL)
- Attribute Importance (AI)

* 11g R1 のみ

Oracle との統合の要件

Oracle Data Mining を使用してデータベース内のモデル作成を実行する場合、以下の条件が前提条件となります。場合によっては、これらの条件が満たされているかデータベース管理者に問い合わせを確認してください。

- ローカル・モードで動作している IBM スポス モデラー、または Windows か UNIX 上の Clementine Server インストール済み環境で動作している IBM スポスモデラー・サーバー。
- Oracle Data Mining オプションを搭載した Oracle 10gR2 または 11gR1 (10.2 データベース 以降)。

注：10gR2 は、一般化線型モデル (11gR1 が必要) 以外のデータベース モデル作成アルゴリズムをサポートします。

- 次に説明するように、Oracle への接続に使用する ODBC データ・ソース。

注：DB モデリングおよび SQL 最適化では、IBM スポスモデラー・サーバー 接続が IBM スポス モデラー コンピュータ上で有効になっている必要があります。この設定を有効にすると、データベース・アルゴリズムにアクセスし、IBM スポス モデラー から SQL を直接プッシュバック、IBM スポスモデラー・サーバーにアクセスできます。現在のライセンスステータスを確認するには、IBM スポス モデラー メニューから次を選択します。

ヘルプ > 製品情報 > 追加の詳細

接続が有効な場合、「ライセンス ステータス」タブにオプション「サーバーの有効化」が表示されます。

Oracle との統合を有効にする

Oracle Data Mining との IBM スポス モデラー の統合を有効にするには、Oracle を設定し、ODBC ソースを作成して、IBM スポス モデラー の「ヘルパー アプリケーション」ダイアログ・ボックスで統合を有効にする必要があります。さらに、SQL の生成と最適化を有効にします。

Oracle の設定

Oracle Data Mining をインストールして構成するには、Oracle の資料 (特に「*Oracle Administrator's Guide*」) を参照してください。

Oracle の ODBC ソースの作成

Oracle と IBM スポス モデラー の接続を有効にするには、ODBC システム・データ・ソース名 (DSN) を作成する必要があります。

DSN を作成する前に、ODBC データ・ソースおよびドライバの基本的な理解し、IBM スポス モデラー のデータベースをサポートする必要があります。

IBM スポスモデラー・サーバー に対して、分散モードで実行している場合、サーバ コンピューターに DSN を作成します。ローカル (クライアント) モードで実行している場合、クライアント・コンピューターに DSN を作成します。

1. ODBC ドライバーをインストールします。このリリースに付属する IBM スポス データ・アクセス・パック インストール・ディスクにあります。*setup.exe* ファイルを実行してインストーラーを起動し、関連するドライバーをすべて選択します。画面上の指示に従って、ドライバーをインストールします。

a. DSN を作成します。

注：メニュー・シーケンスは使用する Windows のバージョンによって異なります。

- **Windows XP の場合。**「スタート」メニューから、「コントロール パネル」を選択します。「管理ツール」をダブルクリックし、次に「データ ソース (ODBC)」をダブルクリックします。
- **Windows Vista。**「スタート」メニューから、コントロール・パネルを選択し、次にシステム保護を選択します。「管理ツール」をダブルクリックし、次に「データ ソース」を選択して「開く」をクリックします。
- **Windows 7。**「スタート」メニューから、コントロール・パネル、システム & セキュリティ、管理ツールの順に選択します。「データ ソース (ODBC)」を選択して「開く」をクリックします。

b. 「システム DSN」 タブに移動してから、「追加」をクリックします。

2. **SPSS OEM 6.0 Oracle Wire Protocol** ドライバーを選択します。
3. 「完了 (Finish)」をクリックします。
4. 「ODBC Oracle Wire Protocol Driver セットアップ」画面では、選択したデータ・ソースの名前、Oracle サーバのホスト名、接続用のポート番号、および使用する Oracle インスタンスの SID を入力します。
tnsnames.ora ファイルで TNS を実装した場合、ホスト名、ポート、および SID は、サーバマシンにある *tnsnames.ora* ファイルから取得できます。詳細は Oracle 管理者までお問い合わせください。
5. 接続をテストするには、「テスト」ボタンをクリックします。

IBM スポス モデラー で Oracle Data Mining Integration の有効化

1. IBM スポス モデラー メニューから次の項目を選択します。

ツール > オプション > ヘルパー・アプリケーション

2. 「Oracle」タブをクリックします。

Oracle Data Mining 統合を有効化: IBM スポス モデラー ウィンドウの下部にデータベース・モデリング・パレットを表示し、Oracle Data Mining アルゴリズムのノードを追加します。

Oracle Connection。 モデルの作成と保管に使用するデフォルトの Oracle ODBC データ・ソースと、有効なユーザー名およびパスワードを指定します。この設定を個々のモデル作成ノードおよびモデル・ナゲットでオーバーライドすることができます。

注: モデル作成の目的で使用するデータベース接続は、データアクセスに使用する接続と同じであっても、異なっても問題ありません。例えば、特定の Oracle データベースのデータにアクセスし、クリーニングまたはその他の操作のためにデータを IBM スポス モデラー にダウンロードし、次にモデル作成を目的としてデータを別の Oracle データベースにアップロードするストリームを利用できます。あるいは、元のデータをフラット・ファイルや他の (Oracle 以外の) ソースに格納できます。この場合、モデルを作成する際に、データを Oracle にアップロードする必要があります。いずれの場合でも、データはモデル作成用のデータベース内で作成された一時テーブルに自動的にアップロードされます。

Oracle データ・マイニング・モデルを上書きしようとしているときに警告します。 データベースに保管されているモデルが警告なしで IBM スポス モデラー によって上書きされないようにするには、このオプションを選択します。

Oracle データ・マイニング・モデルをリストします。 使用可能なデータ・マイニング・モデルを表示します。

Oracle Data Miner の起動を有効にします。(オプション) 有効にすると、IBM スポス モデラー が Oracle Data Miner アプリケーションを起動できるようになります。詳しくは、[42 ページの『Oracle Data Miner』](#)を参照してください。

Oracle Data Miner 実行可能ファイルのパス:(オプション) Windows 版 Oracle Data Miner の実行可能ファイルの物理的な位置を指定します (C:\#xa5;odm\#xa5;bin\#xa5;odminerw.exe など)。Oracle Data Miner は、IBM スポス モデラー とともにインストールされません。Oracle の Web サイト (<http://www.oracle.com/technology/products/bi/odm/odminer.html>) から正しいバージョンをダウンロードし、クライアントでインストールする必要があります。

SQL の生成と最適化を有効にする

1. IBM スポス モデラー メニューから次の項目を選択します。

ツール > ストリームのプロパティ > オプション

2. ナビゲーション ペインの「最適化」 オプションをクリックします。
3. 「SQL 生成」 オプションが有効になっていること確認します。この設定は、データベースのモデル作成が機能するために必要です。
4. 「SQL 生成の最適化」と「その他の実行を最適化」を選択します (絶対に必要な訳ではありませんが、最適化されたパフォーマンスを得るために、選択することを強くお勧めします)。

Oracle Data Mining を使用してモデルを構築する

多少の例外はありますが、Oracle モデル作成ノードは、IBM スポス モデラー の他のモデル作成ノードと同様に動作します。このノードは IBM スポス モデラー ウィンドウの下部にある「モデル作成」パレットからアクセスできます。

データの考慮事項

Oracle では、カテゴリ・データを文字列形式 (CHAR または VARCHAR2) で格納する必要があります。したがって、IBM スポス モデラー では、記憶域のフィールドを、ODM モデルの入力値として「フラグ型」または「名義型」(カテゴリ) が測定の尺度の数値のストレージ・フィールドを指定することは許されません。IBM スポス モデラー でデータ分類ノードを使用すると、必要に応じて、数値を文字列に変換できます。

対象フィールド: ODM 分類モデルでは、出力 (対象) フィールドとして、1 つのフィールドのみを選択できます。

モデル名。 Oracle 11gR1 以降、名前 unique はキーワードで、カスタム・モデル名として使用できません。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

一般的なコメント

- Oracle Data Mining で生成するモデルの場合、IBM スポス モデラー からは PMML のエクスポート/インポートは提供しません。
- モデル・スコアリングは、常に ODM で実行されます。データが IBM スポス モデラー 内にあったり、Clem 内でデータを準備する必要があったりする場合は、データセットを一時テーブルにアップロードする必要がある場合があります。
- IBM スポス モデラー では、通常、単一の予測および関連付けられている確率または確信度のみが提供されます。
- IBM スポス モデラー では、モデル作成とスコアリングに使用できるフィールド数を 1,000 に制限しています。
- IBM スポスモデラー・ソリューション・パブリッシャ を使用して実行用に公開したストリームでは、IBM スポス モデラー で ODM モデルをスコアリングできます。

Oracle モデルの「サーバー」オプション

モデル作成用データのアップロードに使用する Oracle 接続を指定します。「ヘルパー アプリケーション」ダイアログ・ボックスで指定されたデフォルトの Oracle 接続を上書きするために、必要に応じて、「サーバー」タブで各モデル作成ノードに対して接続を選択できます。詳しくは、トピック [26 ページの『Oracle との統合を有効にする』](#)を参照してください。

コメント

- モデル作成に使用する接続は、ストリームの入力ノードで使用する接続と同じであっても別であってもかまいません。例えば、特定の Oracle データベースのデータにアクセスし、クリーニングまたはその他の操作のためにデータを IBM スポス モデラー にダウンロードし、次にモデル作成を目的としてデータを別の Oracle データベースにアップロードするストリームを利用できます。
- ODBC データ・ソース名は、各 IBM スポス モデラー ストリームに効果的に埋め込まれます。あるホスト上で作成されたストリームが別のホスト上で実行された場合、データ・ソースの名前はそれぞれのホストで同じである必要があります。また、各入力ノードまたはモデル作成ノードで、「サーバー」タブから異なるデータ・ソースを選択できます。

誤分類コスト

状況によっては、特定の誤りコストが他の誤りコストに比べて高いことがあります。例えば、信用リスクの高い申請者を低リスクに分類した場合 (ある種の誤分類) のコストは、リスクの低い申請者を高リスクに分類した場合 (別種の誤分類) よりも高くなります。誤分類コストでは、さまざまな予測の誤りに対し、相対的な重要度を指定できます。

誤分類コストは、基本的には、特定の結果に対して適用される重みです。これらの重みは、モデルに組み込まれ、(コストの高い誤りを防ぐための手段として) 実際に予測値に影響する場合があります。

C5.0 モデルを例外として、誤分類コストは、モデルのスコアリング時には適用されず、自動分類ノード、評価グラフ、または精度分析ノードを使用してモデルをランク付けまたは比較する場合には考慮されません。コストを含むモデルは、コストを含まないモデルに比べてエラーが少なく、全体の精度の項目で高くランク付けされません。ただし、コストが少ないエラーにより組み込まれたバイアスがあるため、実際の問題でパフォーマンスが優れる場合があります。

コスト行列には、可能な各予測カテゴリーや実際のカテゴリーの組み合わせのコストが表示されます。デフォルトでは、すべての誤分類コストが 1.0 に設定されています。コストの値を自分で入力するには、「**誤分類コストを使用**」を選択して、コスト行列に独自の値を入力します。

誤分類コストを変更するには、目的の予測値と実際の値の組み合わせに対応するセルを選択して、セルの内容を削除してから、適切なコストを入力してください。コストは自動的に対称的にはなりません。例えば A を B として誤分類した場合のコストを 2.0 に設定しても、B を A として誤分類した場合のコストは、変更しない限りデフォルト値 (1.0) のまま変わりません。

注: ディシジョン・ツリー・モデルのみで構築時のコストを指定できます。

Oracle Naive Bayes

Naive Bayes は、分類の問題に対応する有名なアルゴリズムです。提示されたすべての予測変数は相互に依存関係がないものとして処理されるので、モデルは *naïve* と命名されます。Naive Bayes は拡張性のある高速のアルゴリズムであり、複数の属性と対象属性の組み合わせに対して、条件に関連する確率を計算します。学習データから、個別の確率が計算されます。各入力変数の各値カテゴリーを計算単位とすると、この確率は各対象クラスの確率を表します。

- モデル作成に使用したデータで、モデルの精度をテストする場合、交差検証が使用されます。この検証方法は、モデル作成に利用できるケース数が少ないときに特に便利です。
- モデル出力は、行列形式で参照できます。行列の数値は、条件に関連する確率であり、予測クラス (列) と予測変数値の組み合わせ (行) を関連付けています。

Naive Bayes の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

データ区分データを使用。 データ区分フィールドが定義されている場合、このオプションでは学習用データ区分からのデータのみがモデル構築に使用されるようにします。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『Oracle Data Mining Concepts』を参照してください。

Naive Bayes の「エキスパート」オプション

モデルの作成時に、特定の値または値のペアが学習データ内に十分でない場合、個々の予測属性値または値のペアは無視されます。値を無視する閾値は、学習データ内のレコード数に基づいて分数で指定します。これらの閾値を調整すると、ノイズを減らして、他のデータセットに一般化するモデルの能力を改善できます。

- **シングルトンしきい値。** 特定の予測属性値のしきい値を指定します。特定の値の発生回数が指定した分数以上にならない場合、この値は無視されます。
- **ペアごとのしきい値。** 指定された属性と予測値のペアのしきい値を指定します。特定の値ペアの発生回数は、指定した分数以上になる必要があります。下回った場合、ペアは無視されます。

予測確率。 モデルに、対象フィールドの考えられる結果に対する正しい予測の確率を含めることができます。この機能を有効にするには、「選択」を選択して「指定」ボタンをクリックし、可能性のある出力のいずれかを選択してから、「挿入」をクリックします。

予測セットを使用します。 対象フィールドの考えられるすべての結果について、考えられるすべての結果のテーブルを生成します。

Oracle Adaptive Bayes

Adaptive Bayes Network (ABN) は、Minimum Description Length (MDL) と自動特徴量選択を使用して、Bayesian Network 分類子を構築します。ABN は、Naive Bayes が適さないような環境でも適切に機能しますし、パフォーマンスが低下する可能性はあるもの、他の多くの環境でも少なくとも同程度の機能を発揮します。ABN アルゴリズムでは、簡素化されたディシジョン・ツリー (単一機能)、剪定された Naive Bayes、ブーストされた複数機能モデルという、高度な 3 種類のベイジアンベース・モデルを作成できます。

注: Oracle Adaptive Bayes アルゴリズムは Oracle 12C で削除されたため、Oracle 12C を使用する場合は IBM スポス モデラー ではサポートされていません。 http://docs.oracle.com/database/121/DMPRG/release_changes.htm#DMPRG726 を参照してください。

生成されたモデル

単一機能の作成モードの場合、ABN は人が判読できるルール・セットに基づいて、簡素化されたディシジョン・ツリーを作成します。ビジネス・ユーザーまたはアナリストは、このルール・セットを参照して、モデル予測の基本原則を理解できます。また、適切に対応することや、第三者に論理的に説明することができます。この特長は、Naive Bayes や複数機能モデルに対して、大きな利点になっています。IBM スポス モデラー では、標準的なルール・セットと同様に、これらのルールを参照できます。単純なルール・セットは、次のようになります。

```
IF MARITAL_STATUS = "Married"
AND EDUCATION_NUM = "13-16"
THEN CHURN= "TRUE"
Confidence = .78, Support = 570 cases
```

剪定された Naive Bayes と複数機能モデルは、IBM スポス モデラー では参照できません。

Adaptive Bayes の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

データ区分データを使用。 データ区分フィールドが定義されている場合、このオプションでは学習用データ区分からのデータのみがモデル構築に使用されるようにします。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

モデル タイプ

モデル作成のために、3 種類のモードから適切なモードを選択できます。

- **マルチフィーチャー。** NB モデルや単一機能および複数機能の製品確率モデルなど、多数のモデルを作成して比較します。単一機能モデルは、最も負担が大きいモードであり、通常、結果を出すまでに最も時間がかかります。ルールが生成されるのは、単一機能モデルが最善であると判断された場合のみです。複数機能または NB モデルが選択されると、ルールは生成されません。
- **単一機能。** 一連のルールに基づいて簡素化されたディシジョン・ツリーを作成します。各ルールには、条件とそれぞれの結果に関連付けられた確率が指定されています。ルールは相互に排他的であり、人が判読できる形式で提供されます。この特長は、Naive Bayes や複数機能モデルに対して大きな利点になっています。
- **ナイーブ・ベイズ** 単一の NB モデルを作成し、それを前のグローバル・サンプル (グローバル・サンプル内のターゲット値の分布) と比較します。NB モデルが出力として生成されるのは、事前のグローバル・サンプルよりも、NB モデルの方が、対象値の予測に適していると判断された場合に限られます。それ以外の場合は、モデルは出力として生成されません。

Adaptive Bayes の「エキスパート」オプション

実行時間を制限します。 最大ビルド時間を分単位で指定するには、このオプションを選択します。このオプションを選択すると、作成したモデルの精度が落ちる可能性がありますが、モデルの作成時間を減らすことができます。アルゴリズムは、モデル作成プロセスの各確認ポイントで、指定された時間内に次の確認ポイントまで完了できるかどうかをチェックし、プロセス続行の可否を判断します。そして、指定された制限を越える場合は、利用できる最適なモデルを返します。

最大予測値。 このオプションを使用すると、使用する予測値の数を制限することにより、モデルの複雑さを制限し、パフォーマンスを向上させることができます。予測値は、MDL 法でランク付けされます。この方法は、モデルに含まれている確率を測定する手段として、対象に対する相関関係を示します。

Max Naive Bayes Predictors。 このオプションは、単純ベイズ・モデルで使用される予測値の最大数を指定します。

Oracle Support Vector Machine (SVM)

サポート・ベクター・マシン (SVM) は、マシン学習理論を使用して、データをオーバーフィットすることなしに最高の予測精度を得られるようにする、分類および回帰アルゴリズムです。SVM は、オプションの学習用データの非線形変換を使用し、次にその変換されたデータ中で回帰式を検索して、クラスの分離 (カテゴリ対象の場合)、または対象¹のフィッティング (連続型対象の場合) を行います。Oracle が実装した SVM では、2 つの使用可能なカーネル、線型またはガウスのいずれかを使用して、モデルを作成できます。線型カーネルは、非線型変換をまったく行わないので、生成されたモデルは、本質的に線型回帰モデルです。

詳細は、『Oracle Data Mining Application Developer's Guide』および『Oracle Data Mining Concepts』を参照してください。

Oracle SVM の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『Oracle Data Mining Concepts』を参照してください。

アクティブな学習。 大規模なビルド・セットを処理する方法を提供します。アクティブな学習により、アルゴリズムは、完全な学習データセットに適用する前に、小さなサンプルに基づいて初期モデルを作成し、その結果に基づいてサンプルとモデルを徐々に更新します。モデルが学習データに集束するまで、またはサポート・ベクトルの最大許容数に達するまで、このサイクルは繰り返されます。

カーネル関数: 「線型」または「ガウス」を選択、またはデフォルト「決定システム」のままで、システムが最も適切なカーネルを選択できるようにします。ガウス・カーネルは、より複雑な関係を学習できますが、一般に計算する時間も長くなります。まず、線型カーネルを使用して、もし線型カーネルが良い適合度を発見できなかった場合にのみガウス・カーネルを使用することもできます。これは、もっぱら回帰モデルに当てはまります。回帰モデルではカーネルの選択がより重要になるからです。また、ガウス・モデルを使用して構築された SVM は IBM スポス モデラー で表示できないことにも注意してください。線型カーネルを使用して構築されたモデルは、標準の回帰モデルと同じ方法で IBM スポス モデラー でブラウズできます。

正規化: 連続型入力フィールドと対象フィールドの正規化方法を指定します。「Z-スコア」、「最小-最大」、または「なし」を選択できます。「データの自動準備」チェック・ボックスをオンにすると、正規化が自動的に実行されます。このチェック・ボックスをオフにすると、正規化は手動で行われます。

Oracle SVM の「エキスパート」オプション

カーネル・キャッシュ・サイズ。 ビルド操作中に計算されたカーネルを保管するために使用されるキャッシュのサイズをバイト単位で指定します。通常、キャッシュ・サイズを増やすと構築速度が速くなります。デフォルト値は 50MB です。

収束判定値。 モデル構築の終了前に許容される許容範囲値を指定します。値の範囲は 0 から 1 までです。デフォルト値は 0.001 です。より大きな値では、構築速度が速くなり、またモデルの精度が低下する傾向があります。

標準偏差を指定します。 ガウス・カーネルによって使用される標準偏差パラメーターを指定します。このパラメーターは、モデルの複雑さと他のデータ・セットへの一般化 (データのオーバーフィットとアンダーフィット) との間のトレードオフに影響します。標準偏差値が高くなるにつれて、アンダーフィッティングになる傾向があります。デフォルトでは、このパラメーターは、学習用データから見積もられます。

イプシロンを指定します。 回帰モデルの場合のみ、イプシロン・インセンシティブ・モデルの構築で許容されるエラーの間隔の値を指定します。つまり、大きなエラー (無視できない) 大きなエラーから、(無視できる) 小さなエラーを区別します。値の範囲は 0 から 1 までです。デフォルトでは、これはトレーニング・データから推定されます。

複雑度係数を指定します。 複雑性ファクタは、(学習用データについて計測される) モデルエラーと、モデルの複雑さの各値間のトレードオフ関係を制御して、データがオーバーフィッティングまたは、アンダーフィッティングになることを回避します。この値を大きくするとは、エラーに、より大きなペナルティーを課しことになり、データをオーバーフィッティングするリスクが増大します。この値を小さくするとは、エラーに、より小さなペナルティーを課しことになり、アンダーフィッティングになりやすくなります。

外れ値の比率を指定します。 トレーニング・データ内の希望する外れ値の比率を指定します。One-Class SVM モデルにのみ有効です。複雑度係数の指定 設定と一緒に使用することはできません。

予測確率。 モデルに、対象フィールドの考えられる結果に対する正しい予測の確率を含めることができます。この機能を有効にするには、「選択」を選択して「指定」ボタンをクリックし、可能性のある出力のいずれかを選択してから、「挿入」をクリックします。

予測セットを使用します。 対象フィールドの考えられるすべての結果について、考えられるすべての結果のテーブルを生成します。

Oracle SVM の「重み」オプション

分類モデルでは、重みを使用して、さまざまな考えられる対象値の関連する重要度を指定することができます。例えば、学習データ内のデータ・ポイントがカテゴリ間で現実的に分布していない場合に役に立ちます。重みを使用すると、モデルを偏らせて、データにうまく表れていないカテゴリの補正を行うことができます。対象値の重みが大きくなると、カテゴリの適切な予測の割合が大きくなります。

重みを設定する方法は 3 つあります。

- **学習データに基づく:** これはデフォルトです。重みは、学習データ中のカテゴリの相対度数に基づいて決定されます。
- **すべてのクラスで同じ:** すべてのカテゴリの重みが $1/k$ として定義されます。ここで、 k は対象のカテゴリの数です。
- **カスタム。** 独自の重みを指定することができます。重みの開始値が、すべてのクラスで同じに設定されます。各カテゴリの重みを、ユーザー定義値に調整することができます。特定のカテゴリの重みを調整するには、そのカテゴリに対応するテーブル中の重みのセルを選択し、セルの内容を削除してから、適切な値を入力してください。

すべての重みの合計は、1.0 です。合計が 1.0 にならない場合、値を自動的に正規化するオプションと警告が表示されます。この自動調整によって、重みの制約を強制しながら、カテゴリ間の比率が維持されます。この調整は、任意の時点で「正規化」ボタンをクリックして行うことができます。すべてのカテゴリで値を均等化するためテーブルをリセットするには、「均等化」ボタンをクリックします。

Oracle 一般化線型モデル (GLM)*

一般化線型モデルは、線型モデルによる限定的な仮説を緩和します。例えば、目標変数に正規分布があるという仮説、目標変数に対する予測の効果が本質的に線形となる仮説などです。一般化線型モデルは、対象の分布が、多項分布またはポアソン分布など、非正規分布となる傾向である場合の予測に適切です。同様に、一般化線型モデルは、予測値と目標変数間の関係またはつながりが非線型であると考えられる場合にも役立ちます。

詳細は、『Oracle Data Mining Application Developer's Guide』および『Oracle Data Mining Concepts』を参照してください。

Oracle GLM の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注： このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『Oracle Data Mining Concepts』を参照してください。

正規化: 連続型入力フィールドと対象フィールドの正規化方法を指定します。「**Z-スコア**」、「**最小-最大**」、または「**なし**」を選択できます。「**データの自動準備**」チェック・ボックスをオンにすると、正規化が自動的に実行されます。このチェック・ボックスをオフにすると、正規化は手動で行われます。

欠損値の処理。 入力データの欠損値の処理方法を指定します。

- 「**平均値または最頻値との置換**」を選択すると、数値型属性の欠損値を平均値と置き換え、カテゴリ型属性の欠損値を最頻値と置き換えます。
- 「**完全なレコードのみ使用**」を使用すると、欠損値を含むレコードは無視されます。

Oracle GLM の「エキスパート」オプション

行の重みを使用します。 隣接するドロップダウン・リストをアクティブにするには、このボックスにチェック・マークを付けます。そこから、行の加重係数を含む列を選択できます。

行の診断をテーブルに保存します。 隣接するテキスト・フィールドをアクティブにするには、このボックスにチェック・マークを付けます。このフィールドには、行レベルの診断を含める表の名前を入力できます。

係数の信頼水準。 対象に対して予測される値が、モデルによって計算された信頼区間内にあることの確実性の度合い (0.0 から 1.0 まで)。確信度の境界は、係数の統計で返されます。

対象の参照カテゴリ: 「**カスタム**」を選択して参照カテゴリとして使用する対象フィールドの値を選択するか、デフォルト値「**自動**」のままにします。

Ridge 回帰: Ridge 回帰とは、変数の相関の度合いが高すぎる状態を補正する手法です。「**自動**」オプションを選択して、アルゴリズムがこの方法の使用を制御できるようにしたり、「**無効**」および「**有効**」オプションを使用して手動で制御できるようにします。手動で頂上回帰を有効にすると、隣接するフィールドに値を入力して、システムのデフォルト値を頂上パラメーターで上書きすることができます。

Ridge 回帰の VIF を作成します。 頂上が線形回帰に使用されているときに分散インフレーション係数 (VIF) 統計を生成する場合は、このボックスにチェック・マークを付けます。

予測確率。 モデルに、対象フィールドの考えられる結果に対する正しい予測の確率を含めることができます。この機能を有効にするには、「**選択**」を選択して「**指定**」ボタンをクリックし、可能性のある出力のいずれかを選択してから、「**挿入**」をクリックします。

予測セットを使用します。 対象フィールドの考えられるすべての結果について、考えられるすべての結果のテーブルを生成します。

Oracle GLM の「重み」オプション

分類モデルでは、重みを使用して、さまざまな考えられる対象値の関連する重要度を指定することができます。例えば、学習データ内のデータ・ポイントがカテゴリ間で現実的に分布していない場合に役に立

ちます。重みを使用すると、モデルを偏らせて、データにうまく表れていないカテゴリの補正を行うことができます。対象値の重みが大きくなると、カテゴリの適切な予測の割合が大きくなります。

重みを設定する方法は3つあります。

- **学習データに基づく:** これはデフォルトです。重みは、学習データ中のカテゴリの相対度数に基づいて決定されます。
- **すべてのクラスで同じ:** すべてのカテゴリの重みが $1/k$ として定義されます。ここで、 k は対象のカテゴリの数です。
- **カスタム。** 独自の重みを指定することができます。重みの開始値が、すべてのクラスで同じに設定されます。各カテゴリの重みを、ユーザー定義値に調整することができます。特定のカテゴリの重みを調整するには、そのカテゴリに対応するテーブル中の重みのセルを選択し、セルの内容を削除してから、適切な値を入力してください。

すべての重みの合計は、1.0 です。合計が 1.0 にならない場合、値を自動的に正規化するオプションと警告が表示されます。この自動調整によって、重みの制約を強制しながら、カテゴリ間の比率が維持されます。この調整は、任意の時点で「**正規化**」ボタンをクリックして行うことができます。すべてのカテゴリで値を均等化するためテーブルをリセットするには、「**均等化**」ボタンをクリックします。

Oracle ディジジョン・ツリー

Oracle Data Mining は、一般的な分類と回帰ツリーのアルゴリズムに基づく古典的なディジジョン・ツリー機能を提供します。ODM ディジジョン・ツリー・モデルは、確信度、サポート、および分割基準を含む、それぞれのノードに関する完全な情報を備えています。それぞれのノードの完全なルールが表示され、さらに、欠損値のあるケースにモデルを適用する場合に代わりに使用すべき各ノードに代理変数属性が適用されます。

ディジジョン・ツリーはあらゆる場合に適用可能であり、利用しやすくわかりやすいため、一般的に使用されています。ディジジョン・ツリーは、各潜在的な入力属性をシフトスルーして、最適な「スプリッター」、つまりダウンストリーム・データ・レコードをより同種の母集団に分割する属性分割点 (AGE > 55 など) を検索します。それぞれの分割ディジジョンのあと、ODM は、ツリー全体を成長させ、レコード、項目、または人々の同種の母集団を表すターミナルの「葉」を作成するというプロセスを繰り返します。ルート・ツリー・ノード (例えば、全人口) から見ると、ディジジョン・ツリーは、IF A, then B ステートメントの人間が理解できるルールを提供します。このようなディジジョン・ツリーは、それぞれのツリー・ノードに関するサポートと確信度も提供します。

Adaptive Bayes Network はそれぞれの予測を説明するのに便利な簡潔でシンプルなルールも提供するのに対し、ディジジョン・ツリーはディジジョンを分割するための Oracle Data Mining ルールを提供します。ディジジョン・ツリーは、最良の顧客、健康的な患者、不正に関連する因子などの詳細なプロファイルを開発するのに役立ちます。

ディジジョン・ツリーのモデル・オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは CustomerID などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『Oracle Data Mining Concepts』を参照してください。

不純度メトリック。 各ノードでデータを分割するための最適なテスト質問を探すために使用するメトリックを指定します。最適なスプリッターと分割値は、ノードのエンティティの対象値の等質性が最も増大する場合に得られます。等質性はメトリックに従って測定されます。サポートされるメトリックは、**gini** と **entropy** です。

ディシジョン・ツリーのエキスパート・オプション

最大の深さ。 作成するツリー・モデルの最大の深さを設定します。

ノード内のレコードの最小パーセンテージ。 ノードごとの最小レコード数のパーセンテージを設定します。

分割のレコードの最小パーセンテージ。 親ノード内のレコードの最小数を、モデルの学習に使用されるレコードの合計数に対する割合 (%) として表します。レコード数がこのパーセンテージ未満の場合、分割は行われません。

ノード内の最小レコード数。 返されるレコードの最小数を設定します。

分割の最小レコード数。 値として表される親ノード内のレコードの最小数を設定します。レコード数がこの値未満の場合、分割は行われません。

ルール ID。 チェック・マークが付けられている場合、特定の分割が行われるツリー内のノードを識別するストリングをモデルに組み込みます。

予測確率。 モデルに、対象フィールドの考えられる結果に対する正しい予測の確率を含めることができます。この機能を有効にするには、「**選択**」を選択して「**指定**」ボタンをクリックし、可能性のある出力のいずれかを選択してから、「**挿入**」をクリックします。

予測セットを使用します。 対象フィールドの考えられるすべての結果について、考えられるすべての結果のテーブルを生成します。

Oracle O-Cluster

Oracle O-Cluster のアルゴリズムは、データ母集団内の自然発生的なグループ化を識別します。直交データ区分クラスタリング (O-Cluster) は、階層グリッドベースのクラスタリング・モデルを作成する Oracle 独自のクラスタリング・アルゴリズムであり、すなわち、入力属性空間に軸並行 (直交) データ区分を作成します。このアルゴリズムは回帰的に動作します。結果として得られる階層構造は、属性空間がモザイク風のクラスターのように見える不規則なグリッドを表します。

O-Cluster アルゴリズムは数値属性とカテゴリー属性の両方を取り扱い、ODM は最適なクラスター定義を自動的に選択します。ODM はクラスター詳細情報、クラスター・ルール、クラスター重心値を提供し、母集団をスコアリングするのに使用できます。

O-Cluster の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『*Oracle Data Mining Concepts*』を参照してください。

最大クラスター数。 生成されるクラスターの最大数を設定します。

O-Cluster の「エキスパート」オプション

最大バッファ。 最大バッファ・サイズを設定します。

感度。 新規クラスターを分離するために必要なピーク密度を指定する小数部を設定します。この割合はグローバルな均一濃度に関連付けられます。

Oracle K-Means

Oracle K-means のアルゴリズムは、データ母集団内の自然発生的なグループ化を識別します。K-means のアルゴリズムは距離に基づくクラスタリング・アルゴリズムであり、データをあらかじめ決められた数のクラスターに分割します (ただし、明確に区別できるケースが十分であることを前提とします)。距離に基づくアルゴリズムは距離メトリック (関数) に依存してデータ・ポイント間の類似性を測定します。データ・ポイントは、使用する距離メトリックに従って最も近いクラスターに割り当てられます。ODM は、K-means の拡張バージョンを提供します。

K-means は階層クラスターをサポートし、数値属性とカテゴリー属性を取り扱い、母集団をユーザー指定のクラスター数に分割します。ODM はクラスター詳細情報、クラスター・ルール、クラスター重心値を提供し、母集団をスコアリングするのに使用できます。

K-Means の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『*Oracle Data Mining Concepts*』を参照してください。

クラスター数: 生成されるクラスターの数を設定します。

距離関数。 K-Means クラスタリングに使用する距離関数を指定します。

分割基準。 K-Means クラスタリングに使用する分割基準を指定します。

正規化: 連続型入力フィールドと対象フィールドの正規化方法を指定します。「**Z-スコア**」、「**最小-最大**」、または「**なし**」を選択できます。

K-means の「エキスパート」オプション

反復。 K-means アルゴリズムの反復数を設定します。

収束判定値。 K-Means アルゴリズムの収束判定値を設定します。

ビン数: K-means で作成される属性ヒストグラムにおけるビン数を指定します。それぞれの属性のビンの境界は、学習データセット全体において包括的に計算されます。データ分割手段は固定幅です。ビンが一つしかない単一値の属性を除き、属性はすべて同じ数のビンを備えています。

成長をブロックします。 クラスター・データを保持するために割り振られるメモリーの増加係数を設定します。

最小パーセント属性サポート。 属性がクラスターのルール記述に組み込まれるために非ヌルでなければならない属性値の小数部を設定します。欠損値のあるデータに設定したパラメーター値が高すぎると、ルールが非常に短くなったりあるいは空になることさえあります。

Oracle 非負数マトリックス因数分解 (NMF)

負でないマトリックス因数分解 (NMF) は、大きなデータ・セットを代表的な属性に減らすのに役立ちます。概念的には主成分分析 (PCA) と似ていますが、より多くの属性を処理でき、加法表現モデルでは、NMF は、さまざまなユース・ケースに使用できる強力な最先端のデータ・マイニング・アルゴリズムです。

NMF は、用例テキスト・データなど、多数のデータをより小さく簡潔な表現に縮小するのに使用することが可能であり、データの規模を低減できます (同じ情報量ならば、はるかに少ない変数で保存できます)。

NMF モデルの出力は、SVM などの教師あり学習手法またはクラスタリング技法などの教師なし学習手法を

用いて分析できます。Oracle Data Mining は、NMF および SVM のアルゴリズムを使用して非構造テキスト・データをマイニングします。

NMF の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド。各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注：このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。(11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『*Oracle Data Mining Concepts*』を参照してください。

正規化：連続型入力フィールドと対象フィールドの正規化方法を指定します。「**Z-スコア**」、「**最小-最大**」、または「**なし**」を選択できます。「**データの自動準備**」チェック・ボックスをオンにすると、正規化が自動的に実行されます。このチェック・ボックスをオフにすると、正規化は手動で行われます。

NMF の「エキスパート」オプション

フィーチャーの数を指定します。抽出するフィーチャーの数を指定します。

ランダム シード。NMF アルゴリズムのランダム シードを設定します。

反復数。NMF アルゴリズムの反復数を設定します。

収束判定値。NMF アルゴリズムの収束判定値を設定します。

すべての機能を表示します。最良の機能のみの値ではなく、すべての機能の機能 ID と信頼性を表示します。

Oracle Apriori

Apriori アルゴリズムは、データのアソシエーション・ルールを発見します。例えば、「顧客が剃髪を購入した後、剃ったら、80% の信頼度で、シェービングクリームを購入する」というもの。アソシエーション・マイニングの問題は、2 種類の下位問題に分解できます。

- 最小範囲よりも大きな範囲を持つ、多頻度アイテムセットと呼ばれる項目のすべての組み合わせを検索します。
- 多頻度アイテムセットを使用して希望のルールを生成します。これは、例えば ABC と BC が多頻度である場合、「A は BC を意味する」というルールは、 $\text{support}(BC)$ に対する $\text{support}(ABC)$ の比率が少なくとも最小確信度と同じである場合に、有効であるという考えに基づきます。ABCD が多頻度であるため、このルールは最小範囲を備えることになります。ODM アソシエーションは、結果として生じる単一のルールのみをサポートします (ABC は D を意味します)。

多頻度アイテムセットの数は、最小範囲パラメーターで管理します。作成されるルールの数は、多頻度アイテムセット数と確信度パラメーターで管理します。確信度パラメーターの設定が高すぎる場合、ルールではなくアソシエーション・モデルに多頻度アイテムセットがあるかもしれません。

ODM は Apriori アルゴリズムの SQL ベースの実装です。候補の作成とサポート カウントのステップは、SQL クエリーを用いて実装されます。特化されたインメモリー・データ構造は使用しません。SQL クエリーは、データベース・サーバーで効率的に動作するようにさまざまなヒントで微調整されます。

Apriori の「フィールド」オプション

すべてのモデル作成ノードには、「フィールド」タブがあり、そこからモデルの作成に使用するフィールドを指定できます。

Apriori モデルを作成する前に、アソシエーション・モデル作成の対象項目として使用するフィールドを指定する必要があります。

データ型ノードの設定を使用: このオプションは、ノードに対し、上流のデータ型ノードのフィールド情報を使用するように指示します。これはデフォルトです。

ユーザー設定を使用: このオプションは、ノードに対し、上流のデータ型ノードのフィールド情報ではなく、ここで指定するフィールド情報を使用するように指示します。このオプションを選択した後、トランザクション形式を使用しているかどうかによって、ダイアログ内の残りのフィールドを指定します。

トランザクション形式を使用しない場合、次を指定します。

- **入力:** 入力フィールドを選択してください。これは、データ型ノードのフィールドの役割を「入力」に設定するのと似ています。
- **データ区分:** このフィールドでは、モデル構築の学習、テスト、および検証の各ステージ用に、データを独立したサブセット (サンプル) に分割するフィールドを指定できます。

トランザクション形式を使用する場合、次を指定します。

トランザクション形式を使用します。 アイテムごとの行からケースごとの行にデータを変換する場合は、このオプションを使用します。

このオプションを選択すると、このダイアログ・ボックスの下部にあるフィールド・コントロールを次のように変更します。

トランザクション形式の場合、次を指定します。

- **ID.** リストから ID フィールドを選択します。ID フィールドとして使用できるのは、数値またはシンボル値のフィールドです。選択したフィールドでは、一意の値がそれぞれ、ある分析ユニットを示している必要があります。例えば、マーケット・バスケット分析なら、各 ID が 1 人の顧客を表します。Web ログ分析なら、各 ID が 1 台のコンピューター (IP アドレス) あるいは 1 人のユーザー (ログイン・データ) を表します。
- **内容.** モデルの内容フィールドを指定します。このフィールドには、アソシエーション・モデル作成で関心の対象となる項目が含まれています。
- **データ区分:** このフィールドでは、モデル構築の学習、テスト、および検証の各ステージ用に、データを独立したサブセット (サンプル) に分割するフィールドを指定できます。1 組のサンプルをモデルの生成に使用し、別のサンプルで生成したモデルをテストすることにより、そのモデルが、このデータに似た性質を持つより大きなデータセットにどの程度適用できるかについての良い目安を得ることができます。データ型ノードまたはデータ区分ノードを使用して複数のデータ区分フィールドが定義されている場合、データ区分を使用する各モデル作成ノードの「フィールド」タブで単一のデータ区分フィールドを選択する必要があります (1 つのデータ区分だけが存在している場合、データ分割を有効にすると、そのデータ区分が必ず自動的に使用されます。) また、選択したデータ区分を分析に適用するには、そのノードの「モデル・オプション」タブでデータ区分が有効になっている必要があります (このオプションの選択を解除すると、フィールド設定を変更することなくデータ区分を無効にすることができます)。

Apriori の「モデル」オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド. 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラー には制限があり、このキー・フィールドは数値である必要があります。

注: このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備. (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『*Oracle Data Mining Concepts*』を参照してください。

ルール の 最大長。 ルールの前提条件の最大数を 2 から 20 までの整数で設定します。これにより、ルールの複雑さを制限します。ルールが複雑すぎるまたは特有である場合、またはルール・セットが長すぎて学習できない場合、この設定の数値を減らしてください。

最小信頼性。 最小信頼性レベル (0 から 1 までの値) を設定します。指定された基準より確信度の低いルールは破棄されます。

最小サポート。 最小サポートしきい値 (0 から 1 までの値) を設定します。Apriori は、頻度が最小サポートしきい値を超えるパターンを検出します。

Oracle 最小記述長 (MDL)

Oracle 最小記述長 (MDL) アルゴリズムは、対象属性に最も大きな影響力を持つ属性を識別するのに役立ちます。多くの場合、最も影響力のある属性を知ることが事業をよく理解して管理するのに役立ち、モデル作成作業の簡素化を促進します。さらに、このような属性はモデルを拡張するために追加するデータのタイプを示します。MDL は、製造部品の品質、顧客離れに関連の因子、または特定の疾患の治療に最も関連していると思われる遺伝子を予測するのに最適なプロセス属性を検索するのに使用できます。

Oracle MDL は、対象を予測する場合に重要でないものと見なす入力フィールドを破棄します。残りのフィールドで、Oracle Data Miner で表示可能な、Oracle モデルに関連する未調整モデル・ナゲットを作成します。Oracle Data Miner でモデルを参照すると、残りの入力フィールドを、対象を予測する際の重要度の順に示すグラフが表示されます。

負の順位は、ノイズを示します。0 以下にランクされている入力フィールドは予測に貢献せず、データから削除する必要があります。

グラフを表示するには

1. 「モデル」パレットにある未調整モデル・ナゲットを右クリックして、「参照」を選択します。
2. モデル・ウィンドウで、ボタンをクリックして Oracle Data Miner を起動します。
3. Oracle Data Miner に接続します。詳しくは、トピック [42 ページの『Oracle Data Miner』](#) を参照してください。
4. Oracle Data Miner のナビゲータ・パネルで「モデル」、「属性の重要度」を展開します。
5. 関連する Oracle モデルを選択します (IBM スポス モデラーで指定した対象フィールドと同じ名前です)。どれが適切なモデルがわからない場合、Attribute Importance フォルダーを選択して、作成日によってモデルを検索します。

MDL のモデル・オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

固有フィールド。 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。IBM スポス モデラーには制限があり、このキー・フィールドは数値である必要があります。

注： このフィールドは、Oracle Adaptive Bayes、Oracle O-Cluster および Oracle Apriori を除くすべての Oracle ノードについてオプションです。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『Oracle Data Mining Concepts』を参照してください。

Oracle Attribute Importance (AI)

属性の重要度の目的は、結果に関連するデータセットの属性、および最終的な結果に影響を与える程度を見出すことです。Oracle Attribute Importance モデル作成ノードは、データを分析、パターンを検索、関連レベルの確信度の結果を予測します。

AI のモデル・オプション

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

データ区分データを使用。 データ区分フィールドが定義されている場合、このオプションでは学習用データ区分からのデータのみがモデル構築に使用されるようにします。

データの自動準備。 (11 のみ) Oracle データ・マイニングの自動データ準備モードを有効 (デフォルト) または無効にします。このボックスがチェックされている場合、ODM は、アルゴリズムに必要なデータ変換を自動的に実行します。詳細は、『Oracle Data Mining Concepts』を参照してください。

AI の選択オプション

「オプション」タブで、モデル・ナゲット内の入力フィールドを選択または除外するデフォルトの設定を指定できます。その後、以後のモデル構築作業で使用するフィールドのサブセットを選択するために、ストリームへモデルを追加できます。または、モデルの生成後にモデル・ブラウザー内で追加のフィールドを選択したり選択を解除したりして、このような設定を上書きすることもできます。ただし、デフォルトの設定はそれ以上変更しなくてもモデル・ナゲットに適用できるので、スクリプトを作成する目的に対しては、特に有用です。

The following options are available:

ランク付けされているすべてのフィールド: 「重要」、「境界」、または「重要ではない」のランクに基づいてフィールドを選択します。各ランクと、レコードにランクを割り当てるために使用される分割値のラベルは、編集できます。

フィールドの上位数: 重要度に基づいて上位 n 件のフィールドを選択します。

次より大きい重要度。 指定された値よりも高い重要度のすべてのフィールドを選択します。

対象フィールドは、この選択にかかわらず、常に保存されます。

AI モデル・ナゲットの「モデル」タブ

Oracle AI モデル・ナゲットの「モデル」タブには、すべての入力フィールドのランクと重要度が表示されるので、フィルタリングするフィールドを、左の列のチェック・ボックスを使用して選択できるようになります。ストリームを実行すると、対象の予測とともに、チェックマークが付けられたフィールドだけが保存されます。その他の入力フィールドは、廃棄されます。デフォルトの選択はモデル作成ノード内で指定されたオプションに基づきますが、必要に応じて追加のフィールドを選択したり、選択を解除したりできます。

- ランク、フィールド名、重要度、またはその他の表示された列でリストをソートするには、列見出しをクリックします。または、ツールバーを使用して、「ソート項目」ボタンの隣のリストから該当する項目を選択し、上方向矢印と下方向矢印を使用してソートの方向を変更します。
- ツールバーを使用してすべてのフィールドにチェックを入れたり外したりできます。また、「フィールドのチェック」ダイアログ・ボックスを利用してランクまたは重要度でフィールドを選択できます。Shift キーまたは Ctrl キーを押してフィールドをクリックすると、複数選択することもできます。
- 重要度が高い、境界、重要度が低い、として入力フィールドをランク付けするための閾値は、テーブルの下の凡例に表示されます。これらの値は、モデル作成ノード内で指定されます。

Oracle モデルの管理

Oracle モデルは、他の IBM スポス モデラー のモデルのように「モデル」パレットに追加され、ほとんど同じように使用できます。ただし、IBM スポス モデラー 中で生成された各 Oracle モデルは実際にはデータベース・サーバー上にあるモデルへの参照であるなどの、重要な違いがいくつかあります。

Oracle モデル・ナゲットの「サーバー」タブ

ODM モデルの構築は、IBM スポス モデラー を通じて IBM スポス モデラー 内でモデルを作成し、さらに、Oracle データベース内でモデルを作成または置換します。この種類の IBM スポス モデラー モデルは、デ

データベース・サーバーに格納されているデータベース・モデルの内容を参照します。IBM スポス モデラーは、IBM スポス モデラー モデルと Oracle モデルの両方に、同一の**モデル・キー**文字列を生成して格納し、整合性チェックを実行します。

それぞれの Oracle モデル用のキー文字列は、「モデルの一覧」ダイアログ・ボックスの「モデル情報」列に表示されます。IBM スポス モデラー モデルのキー文字列は、IBM スポス モデラー モデルの「サーバー」タブの「**モデル・キー**」として表示されます (ストリームに置かれた場合)。

モデル・ナゲットのダイアログ・ボックスにある「**検査**」ボタンは、IBM スポス モデラー モデルと Oracle モデルのキーが一致するかどうかを検査するために使用できます。Oracle に同じ名前のモデルがないかモデル・キーが一致しない場合は、Oracle モデルは削除されているか、または、IBM スポス モデラー モデルの作成によって再構築されています。

Oracle モデル・ナゲットの「要約」タブ

モデル・ナゲットの「要約」タブで、モデルそのもの (精度分析)、モデルで使用するフィールド (フィールド)、モデルの構築時に使用する設定 (構築の設定)、およびモデルの学習 (学習の要約) についての情報を表示します。

ノードを初めて参照する場合、「要約」タブの結果は閉じられています。目的の結果を表示するには、項目の左側にある展開コントロールを使用して項目を展開するか、または「**すべて展開**」ボタンをクリックしてすべての結果を表示します。見終わった結果を隠すには、展開コントロールを使用して目的の結果を省略するか、または「**すべて閉じる**」ボタンをクリックしてすべての結果を非表示にします。

分析。 特定のモデルについての情報を表示します。このモデル・ナゲットに接続されている精度分析ノードを実行した場合、その精度分析情報もこのセクションに表示されます。

フィールド。 対象フィールドおよびモデル構築時の入力として使われるフィールドが表示されます。

構築の設定: モデル構築時に使われる設定情報が表示されます。

学習の要約: モデルの種類、モデルの作成に使われたストリーム、モデルの作成者、モデルの作成日時、およびモデルの構築時間などの情報が表示されます。

Oracle モデル・ナゲットの「設定」タブ

モデル・ナゲットの「設定」タブで、モデル作成ノードの特定のオプションの設定を、スコアリングの目的で優先させることができます。

Oracle ディシジョン・ツリー

誤分類コストを使用します。 Oracle ディシジョン・ツリー・モデルで誤分類コストを使用するかどうかを決定します。詳しくは、トピック [28 ページの『誤分類コスト』](#)を参照してください。

ルール ID。 これを選択 (チェック・マークを付ける) すると、ルール ID 列が Oracle ディシジョン・ツリー・モデルに追加されます。ルール識別子は、特定の分割が作成されるツリーのノードを識別します。

Oracle NMF

すべての機能を表示します。 これを選択した場合 (チェック・マークが付けられている場合)、Oracle NMF モデルでは、最良の特徴量のみの値ではなく、すべての特徴量の特徴量 ID と確信度が表示されます。

Oracle モデルのリスト作成

「Oracle Data Mining Model の一覧」ボタンをクリックすると、既存のデータベース・モデルの一覧を表示するダイアログ・ボックスが表示され、モデルを削除できます。このダイアログ・ボックスは、「ヘルパーアプリケーション」ダイアログ・ボックスおよび ODM 関連のノードの構築、ブラウズ、および適用の各ダイアログ・ボックスから起動できます。

各モデルについて、次の情報が表示されます。

- **モデル名。** リストをソートするのに使用されるモデルの名前
- **モデル情報。** ビルド日時とターゲット列名で構成されるモデル・キー情報
- **モデル・タイプ。** このモデルの構築に使用されたアルゴリズムの名前

Oracle Data Miner

Oracle Data Miner は Oracle Data Mining (ODM) に対するユーザー・インターフェースであり、以前の IBM スポス モデラー のユーザー・インターフェースを ODM 用に置き換えます。Oracle Data Miner は、ODM アルゴリズムの適切な活用においてアナリストの成功率を上げるように設計されています。このような目標は、いくつかの方法で対処します。

- ユーザーは、データの準備とアルゴリズムの選択の両方に対処する手法を適用する場合により多くの支援を必要とします。Oracle Data Miner は、データ・マイニング活動を提供して適切な手法をユーザーに示すことにより、この必要性を満たします。
- Oracle Data Miner は、改良および拡張された発見的手法をモデルの構築と変換のウィザードに含めて、モデルと変換の設定を指定する場合におけるエラー発生機会を低減します。

Oracle Data Miner の接続の定義

1. Oracle Data Miner は、「**Oracle Data Miner の起動**」ボタンを介して、Oracle の「構築」、「ノードの適用」、および「出力」のどのダイアログ・ボックスからでも起動できます。



図 2. 「Oracle Data Miner の起動」ボタン

2. Oracle Data Miner の「**接続の編集**」ダイアログ・ボックスは、Oracle Data Miner の外部アプリケーションが起動する前にユーザーに示されます (ただし、「ヘルパー アプリケーション」オプションが適切に定義されていることが前提です)。

注: このダイアログ・ボックスは、定義された接続名がない場合にのみ表示されます。

- Data Miner の接続名を指定し、適切な Oracle 10gR1 または 10gR2 のサーバー情報を入力します。Oracle サーバーは IBM スポス モデラー で指定されているのと同じサーバーにする必要があります。
3. Oracle Data Miner の「**接続の選択**」ダイアログ・ボックスは、使用する接続名 (上記のステップで定義) を指定するためのオプションを示します。

Oracle Data Miner の要件、インストール、および使用に関する詳細は、Oracle の Web サイトの [Oracle Data Miner](#) を参照してください。

データの準備

Oracle Data Mining により提供される Naive Bayes、Adaptive Bayes、および Support Vector Machine のいずれかを使用してモデルを作成する場合、2 種類のデータを準備すると便利です。

- **データ分割**とは、連続型データを使用できないアルゴリズムのために、連続した数値範囲フィールドをカテゴリへの変換することです。
- **正規化**とは、ある数値の範囲が似かよった平均値と標準偏差を持つようにするために、数値範囲に適用される変換です。

データ分割

IBM スポス モデラー のデータ分割ノードは、データ分割操作を実行するための数多くのテクニックを提供しています。同じデータ分割操作が、単一のフィールドにも複数のフィールドにも適用できるように定義されています。データ・セットに対してデータ分割操作を実行すると、閾値が生成され、IBM スポス モデラー のフィールド作成ノードが作成できるようになります。フィールド作成操作は、SQL に変換でき、モデルの構築およびスコアリングの前に適用できます。このアプローチでは、モデルと、データ分割を実行するフィールド作成ノードの間に依存関係が生じますが、データ分割の指定は、複数のモデル作成タスクで再利用できます。

正規化

Support Vector Machine モデルへの入力として使用される連続型 (数値範囲) フィールドは、モデルを構築する前に正規化する必要があります。回帰のモデルの場合、正規化は、モデル出力からスコアを再構築するためにも必要です。SVM モデルの設定では、「**Z-スコア**」、「**Min-Max**」、または「**なし**」を選択できます。正規化係数は、Oracle によりモデル構築プロセスの 1 ステップとして作成されます。そして、その係数は IBM スポス モデラー にアップロードされ、そのモデルにも格納されます。適用時に、係数は IBM ス

ポス モデラー のフィールド作成式に変換され、スコアリング用のデータをモデルに渡す前に、そのデータの準備に使用されます。この場合、正規化は、モデル作成タスクと密接に関係しています。

Oracle Data Mining の例

IBM スポス モデラー と共に ODM を使用方法について解説する、数多くのサンプル・ストリームが含まれています。これらのストリームは、`¥Demos¥Database_Modelling¥Oracle Data Mining¥` の IBM スポス モデラー インストール・フォルダーにあります。

注: Demos フォルダーには、Windows の「スタート」メニューの IBM スポス モデラー プログラム・グループからアクセスできます。

次の表に示すストリームは、Oracle Data Mining に組み込まれているサポート・ベクター・マシン (SVM) アルゴリズムを使用したデータベース・マイニング処理の例として、まとめて順に使用することができます。

表 4. データベース・マイニング - ストリーム例	
ストリーム	説明
<code>1_upload_data.str</code>	フラット・ファイルのデータを整理して、データベースへアップロードするために使用されます。
<code>2_explore_data.str</code>	IBM スポス モデラー でのデータ探索の例として使用されます。
<code>3_build_model.str</code>	データベース固有のアルゴリズムを使用したモデルを構築します。
<code>4_evaluate_model.str</code>	IBM スポス モデラー でのモデル評価の例として使用されます。
<code>5_deploy_model.str</code>	データベース内スコアリングのためにモデルを展開します。

注: サンプルを実行するには、ストリームを順番に実行する必要があります。さらに、各ストリーム中の入力およびモデル作成ノードは、使用するデータベースの有効なデータ・ソースを参照するように更新する必要があります。

サンプル・ストリームで使用されるデータセットは、クレジット カード申請に関するものであり、カテゴリ型および連続型予測フィールドの混在について、分類上の問題を提示します。このデータ・セットの詳細は、サンプル・ストリームと同じフォルダーにある `crx.names` ファイルを参照してください。

このデータセットは、<ftp://ftp.ics.uci.edu/pub/machine-learning-databases/credit-screening/> にある UCI Machine Learning Repository から入手可能です。

ストリームの例: データのアップロード

最初のストリーム例 `1_upload_data.str` は、フラット・ファイルのデータを整理して Oracle へアップロードするのに使用されます。

Oracle Data Mining には、一意の ID フィールドが必要なので、この初期ストリームは、フィールド作成ノードを使用して、ユニークな値 1、2、3 を持つ `ID` という名前の新しいフィールドをデータセットに追加します。このフィールド作成ノードでは、IBM スポス モデラー の `@INDEX` が使用されています。

置換ノードは、欠損値の処理に使用され、`crx.data` テキスト・ファイルから読み込まれた空のフィールドを `NULL` 値で置き換えます。

ストリームの例: データの調査

2 番目の例のストリーム、`2_explore_data.str` を使用して、要約統計およびグラフなど、データの概要を取得するデータ検査ノードの使用方法を説明します。

データ検査レポート内のグラフをダブルクリックすると、指定されたフィールドについて、より深く検索した結果を表す詳細なグラフが生成されます。

ストリームの例:モデルの作成

3 番目のストリーム例、*3_build_model.str* では、IBM スポス モデラー でのモデル構築を説明します。データベース入力ノード (CREDIT) をダブルクリックして、データ・ソースを指定します。構築設定を指定するには、構築ノード (最初のラベルは CLASS、データ・ソースを指定すると FIELD16 に変更) をダブルクリックします。

ダイアログ・ボックスの「モデル」タブで、次の作業を行います。

1. 一意のフィールドとして「**ID**」を選択します。
2. カーネル 関数として「**線型**」を選択し、正規化の方法として「**Z-スコア**」を選択します。

ストリームの例:モデルの評価

第 4 のストリーム例、*4_evaluate_model.str* では、IBM スポス モデラー を使用したデータベース内モデル作成の利点を説明します。モデルを実行すると、その結果をユーザーのデータストリームに追加したり、IBM スポス モデラー が提供するいくつかのツールを使用してモデルを評価したりできます。

モデル作成結果の表示

テーブル・ノードをモデル・ナゲットに適用して、結果を検証します。**\$O-field16** フィールドに各ケースの *field16* の予測値が表示され、**\$OC-field16** にこの予測の確信度値が表示されます。

モデル作成結果の評価

精度分析ノードを使用すると、予測フィールドとその対象フィールド間の一致パターンを表す一致行列を作成できます。結果を表示するには精度分析ノードを実行します。

評価ノードを使用して、このモデルによりどの程度、精度が改善されたかを示すゲイン・グラフも作成できます。結果を表示するには評価ノードを実行します。

ストリームの例:モデルの展開

モデルの精度が満足できるものになったら、外部のアプリケーションと共に使用するために展開したり、データベースに保存するために発行できます。最後のストリームの例、*5_deploy_model.str* では、データはテーブル CREDITDATA から読み込まれ、スコアリングされてから、*deploy solution* という Publisher ノードを使用して テーブル CREDITSCORES に発行されます。

第 5 章 IBM Data Warehouse および IBM NetezzaAnalytics によるデータベース・モデリング

スpos モデラー と、IBM Data Warehouse および IBM NetezzaAnalytics

IBM スpos モデラー では、IBM Data Warehouse および IBM Netezza[®]Analytics との統合をサポートしています。この統合により、それらの IBM サーバーで高度な分析を実行する機能が提供されます。これらの機能は、IBM スpos モデラー のグラフィカル・ユーザー・インターフェースとワークフロー指向の開発環境で使用することができ、IBM Netezza 環境または IBM Data Warehouse 環境で直接データ・マイニング・アルゴリズムを実行できます。

スpos モデラー は、**IBM NetezzaAnalytics** の次のアルゴリズムの統合をサポートします。

- ディジジョン・ツリー
- K-Means
- TwoStep
- ベイズ・ネット
- Naive Bayes
- KNN
- 分裂クラスタリング
- PCA
- 回帰ツリー
- 線型
- 時系列
- 一般化線型

これらのアルゴリズムの詳細は、「*IBM NetezzaAnalytics* 開発者ガイド」および「*IBM NetezzaAnalytics* リファレンス・ガイド」を参照してください。

スpos モデラー は、**IBM Data Warehouse** の次のアルゴリズムの統合をサポートします (ベイズ・ネット、分裂クラスタリング、および時系列はサポートされません)。

- ディジジョン・ツリー
- K-Means
- TwoStep
- Naive Bayes
- KNN
- PCA
- 回帰ツリー
- 線型
- 一般化線型

注：AIX はサポートされていません。

統合の要件

IBM NetezzaAnalytics または IBM Data Warehouse を使用してデータベース内のモデル作成を実行する場合、以下の条件が前提条件となります。場合によっては、これらの条件が満たされているかデータベース管理者に問い合わせて確認してください。

- Windows または UNIX (IBM Netezza ODBC ドライバーを使用できない zLinux を除く) 上の IBM スポス モデラー・サーバー インストール済み環境に対して稼働する IBM スポス モデラー。
- IBM NetezzaAnalytics パッケージを実行する IBM Netezza Performance Server。

注: 必要な Netezza Performance Server (NPS) の最小バージョンは、必要な INZA のバージョンによって異なり、次のようになります。

- NPS 6.0.0 P8 より大きなバージョンはすべて、2.0 より前の INZA のバージョンをサポートします。
- INZA 2.0 以上を使用するには、NPS 6.0.5 P5 以上が必要です。

Netezza 一般化線形および Netezza 時系列が機能するためには、INZA 2.0 以上が必要です。その他すべての Netezza データベース内ノードには、INZA 1.1 以降が必要です。

- IBM Netezza データベースに接続するための ODBC データソース。詳しくは、トピック [46 ページの『統合の有効化』](#) を参照してください。
- IBM Data Warehouse データベースに接続するための ODBC データ・ソース。
- IBM スポス モデラー で有効化された SQL の生成および最適化。詳しくは、トピック [46 ページの『統合の有効化』](#) を参照してください。

注: DB モデリングおよび SQL 最適化では、IBM スポスモデラー・サーバー 接続が IBM スポス モデラー コンピュータ上で有効になっている必要があります。この設定を有効にすると、データベース・アルゴリズムにアクセスし、IBM スポス モデラー から SQL を直接プッシュバック、IBM スポスモデラー・サーバーにアクセスできます。現在のライセンスステータスを確認するには、IBM スポス モデラー メニューから次を選択します。

ヘルプ > 製品情報 > 追加の詳細

接続が有効な場合、「ライセンス ステータス」タブにオプション「**サーバーの有効化**」が表示されます。

統合の有効化

IBM NetezzaAnalytics または IBM Data Warehouse との統合を有効にする手順は、次のとおりです。

- IBM NetezzaAnalytics または IBM Data Warehouse を構成する
- ODBC ソースの作成
- IBM スポス モデラー で統合を有効にする
- IBM スポス モデラー での SQL の生成と最適化を有効にする

これらについては、後の項で説明します。

IBM NetezzaAnalytics または IBM Data Warehouse を構成する

IBM NetezzaAnalytics または IBM Data Warehouse をインストールして構成するには、該当する IBM 資料を参照してください。例えば、IBM NetezzaAnalytics の場合は、その製品に付属の「*IBM NetezzaAnalytics* インストール・ガイド」を参照してください。ガイドの「データベース権限の設定」に、IBM スポス モデラー ストリームをデータベースに書き込むために実行する必要のあるスクリプトの詳細について記載されています。

注: 行列の計算が必要なノードを使用する場合、CALL NZM..INITIALIZE(); を実行して Matrix Engine を初期化する必要があります。これを初期化しないと、ストアード・プロシージャの実行が失敗します。各データベースごとに一度のセットアップ手順で初期化できます。

IBM NetezzaAnalytics の ODBC ソースの作成

IBM Netezza database と IBM スポス モデラー の接続を有効にするには、ODBC システム・データ・ソース名 (DSN) を作成する必要があります。

DSN を作成する前に、ODBC データ・ソースおよびドライバの基本的な理解し、IBM スポス モデラー のデータベースをサポートする必要があります。

IBM スポスモデラー・サーバー に対して、分散モードで実行している場合、サーバ コンピューターに DSN を作成します。ローカル (クライアント) モードで実行している場合、クライアント・コンピューターに DSN を作成します。

Windows クライアント

1. *Netezza Client* の CD の *nzodbcsetup.exe* ファイルを実行してインストーラーを起動してください。画面上の指示に従って、ドライバをインストールします。詳細は、『IBM Netezza ODBC、JDBC、および OLE DB インストールおよび設定ガイド』を参照してください。

a. DSN を作成します。

注：メニュー・シーケンスは使用する Windows のバージョンによって異なります。

- **Windows XP の場合。** 「スタート」メニューから、「コントロール パネル」を選択します。「管理ツール」をダブルクリックし、次に「データ ソース (ODBC)」をダブルクリックします。
- **Windows Vista。** 「スタート」メニューから、コントロール・パネルを選択し、次にシステム保護を選択します。「管理ツール」をダブルクリックし、次に「データ ソース」を選択して「開く」をクリックします。
- **Windows 7。** 「スタート」メニューから、コントロール・パネル、システム & セキュリティ、管理ツールの順に選択します。「データ ソース (ODBC)」を選択して「開く」をクリックします。

b. 「システム DSN」 タブに移動してから、「追加」をクリックします。

2. リストから「**NetezzaSQL**」を選択し、「完了」をクリックします。
3. 「Netezza ODBC ドライバセットアップ」画面の「**DSN オプション**」タブで、選択したデータ・ソース名、Netezza サーバーのホスト名または IP アドレス、接続のポート番号、使用している IBM Netezza インスタンスのデータベース、データベース接続のためのユーザー名およびパスワードの詳細を入力します。フィールドについての説明は「ヘルプ」ボタンをクリックします。
4. 「**テスト接続**」ボタンをクリックして、データベースに接続できることを確認します。
5. 正常に接続が行われたら、ODBC Data Source Administrator の画面が終了するまで繰り返し「**OK**」を押します。

Windows サーバー

Windows サーバーの手順は、Windows XP のクライアントの手順と同じです。

UNIX または Linux サーバー

以下の手順は、UNIX または LINUX サーバー (IBM Netezza ODBC ドライバを使用できない zLinux を除く) に適用されます。

1. Netezza Client CD/DVD から、関連する <platform>cli.package.tar.gz ファイルをサーバー上の一時的な場所にコピーします。
2. **gunzip** および **untar** コマンドを使用してアーカイブ・コンテンツを抽出します。
3. 実行権限を抽出された *unpack* スクリプトに追加します。
4. 画面上のプロンプトに回答して、スクリプトを実行します。

5. modelersrv.sh ファイルを編集して以下の行を含めます。

```
. <SDAP Install Path>/odbc.sh
LD_LIBRARY_PATH_64=$LD_LIBRARY_PATH:/usr/local/nz/lib64; export LD_LIBRARY_PATH_64
NZ_ODBC_INI_PATH=<SDAP Install Path>; export NZ_ODBC_INI_PATH
```

以下に例を示します。

```
. /usr/IBM/SPSS/SDAP/odbc.sh
LD_LIBRARY_PATH_64=$LD_LIBRARY_PATH:/usr/local/nz/lib64; export LD_LIBRARY_PATH_64
NZ_ODBC_INI_PATH=/usr/IBM/SPSS/SDAP; export NZ_ODBC_INI_PATH
```

6. ファイル /usr/local/nz/lib64/odbc.ini を検索して SDAP とともにインストールされている odbc.ini ファイル (\$ODBCINI 環境変数によって定義されたファイル) にコピーします。

注: 64 ビット Linux システムの場合、**Driver** パラメーターは誤って 32 ビット・ドライバーを参照します。前の手順で odbc.ini コンテンツをコピーする場合、以下のようにパラメーター内のパスを編集します。

```
/usr/local/nz/lib64/libnzodbc.so
```

7. Netezza DSN 定義のパラメーターを編集して使用されるデータベースを反映させます。
8. IBM スポスモデラー・サーバー を再起動して、クライアント上の Netezza データベース内マイニング・ノードの使用をテストします。

スpos モデラー で統合を有効にする

1. IBM スポス モデラー のメイン・メニューから次の各項目を選択します。

ツール > オプション > ヘルパー・アプリケーション。

2. 「**IBM Data Warehouse**」タブをクリックします。

IBM Data Warehouse Analytics 統合を有効化: IBM スポス モデラー ウィンドウの下部でデータベース・モデリング・パレットを有効にし (まだ表示されていない場合)、IBM Data Warehouse と Netezza Data Mining のアルゴリズムのノードを追加します。

IBM Data Warehouse 接続: 「編集」ボタンをクリックし、ODBC ソースの作成時にセットアップした IBM Data Warehouse 接続文字列を選択します。詳しくは、IBM Data Warehouse 管理コンソールを参照してください。

SQL の生成と最適化を有効にする

非常に大きいデータ・セットを扱う確率が高いため、パフォーマンス上の理由により、IBM スポス モデラー で SQL 生成および最適化を有効にする必要があります。

1. IBM スポス モデラー メニューから次の項目を選択します。

ツール > ストリームのプロパティ > オプション

2. ナビゲーション ペインの「**最適化**」オプションをクリックします。
3. 「**SQL 生成**」オプションが有効になっていること確認します。この設定は、データベースのモデル作成が機能するために必要です。
4. 「**SQL 生成の最適化**」と「**その他の実行を最適化**」を選択します (絶対に必要な訳ではありませんが、最適化されたパフォーマンスを得るために、選択することを強くお勧めします)。

IBM NetezzaAnalytics および IBM Data Warehouse でのモデルの構築

サポートされているアルゴリズムのそれぞれに、対応するモデル作成ノードがあります。ノード・パレットの「**DB モデル作成**」タブから、IBM Data Warehouse と IBM Netezza のモデル作成ノードにアクセスできます。

データの考慮事項

データ・ソース内のフィールドには、モデル作成ノードに応じてさまざまなデータ型の変数を含めることができます。IBM スポス モデラー では、データ型は測定の尺度とも呼ばれます。モデル作成ノードの「フィールド」タブでは、入力フィールドと対象フィールドに許可されている測定の尺度の種類を示すアイコンを使用しています。

対象フィールド: 対象フィールドは、値を予測しようとしているフィールドです。対象を指定できる場合、ソース・データ・フィールドのうち 1 つだけを対象フィールドとして選択できます。

レコード ID フィールド: 各ケースを一意に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。ソースデータに ID フィールドが含まれていない場合は、次の手順に示すように、フィールド作成ノードを使用してこのフィールドを作成することができます。

1. 入力ノードを選択します。
2. ノード・パレットの「フィールド設定」タブから、フィールド作成ノードをダブルクリックします。
3. 領域内のアイコンをダブルクリックしてフィールド作成ノードを開きます。
4. 「派生フィールド」フィールドで、例えば ID を入力します。
5. 「**CLEM 式**」フィールドで、@INDEX と入力して「**OK**」をクリックします。
6. フィールド作成ノードをストリームの残りに接続します。

注: NUMERIC(18,0) データ・タイプを使用して Netezza データベースから長い数値データを取得する場合、スポス モデラー はインポート時にデータを切り上げることがあります。この問題を回避するために、データは BIGINT または NUMERIC(36,0) のいずれかのデータ型を使用して保管します。

注: 使用できるフィールド・タイプの制限により、「レコード ID」の役割およびデータ型不明の「測定の尺度」があるフィールドは、Netezza データベース内モデル作成ノード (例えば、K-Means) には表示されません。

ヌル値の処理

入力データに null 値が含まれている場合、いくつかの Netezza ノードを使用することによってエラー・メッセージが表示されたりストリームの実行時間が長くなる可能性があるため、null 値を含むレコードを削除することをお勧めします。以下の方法を使用します。

1. 条件抽出ノードを入力ノードに接続します。
2. 条件抽出ノードの「**モード**」を「**破棄**」に設定します。
3. 「**条件**」フィールドに以下を入力します。

```
@NULL(field1) [or @NULL(field2)[... or @NULL(fieldN)]]
```

すべての入力フィールドを含むようにします。

4. 条件抽出ノードをストリームの残りに接続します。

モデル出力

Data Warehouse または Netezza のモデル作成ノードを含むストリームは、実行のたびにわずかに異なる結果を生成する可能性があります。それは、データがモデル作成の前に一時テーブルに読み込まれるため、ノードがソース・データを読み取る順序が必ずしも同じではないからです。ただし、この効果によって生まれる相違点は無視できます。

一般的なコメント

- IBM スポス コラボレーションおよびデプロイメント・サービス では、IBM Data Warehouse または IBM Netezza のデータベース・モデル作成ノードを含むストリームを使用してスコアリング構成を作成することはできません。
- Data Warehouse または Netezza のノードで作成されたモデルの場合、PMML エクスポートや PMML インポートを行うことはできません。

フィールド・オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定 of 尺度を示します。

「すべて」ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定 of 尺度のボタンをクリックして、その測定 of 尺度のすべてのフィールドを選択します。

target. 1 つのフィールドを予測の対象として選択します。一般化線形モデルの場合、この画面の「試行回数」フィールドも参照してください。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

サーバー・オプション

「サーバー」タブでは、モデルが構築される IBM Data Warehouse データベースを指定します。

IBM Data Warehouse サーバーの詳細: ここで、モデルに使用するデータベースの接続の詳細を指定します。

- **上流の接続を使用:** (デフォルト) データベース入力ノードなど、上流のノードで指定した接続の詳細を使用します。このオプションは、すべての上流ノードが SQL プッシュバックを使用できる場合にのみ有効です。この場合、SQL が完全にすべての上流ノードを実装するため、データベース外のデータを移動する必要はありません。
- **接続へのデータの移動:** ここで指定したデータベースにデータを移動します。データを移動することにより、データが別の IBM Data Warehouse データベース、他のベンダーのデータベースにある場合、またはデータがフラット・ファイルにある場合でもモデリングを実行できるようにします。また、ノードが SQL プッシュバックを実行していないためデータが抽出されていない場合、データはここで指定されたデータベースに戻されます。接続を参照して選択するには、「編集」ボタンをクリックします。



注意: 通常、IBM NetezzaAnalytics と IBM Data Warehouse は、非常に大きいデータ・セットで使用されます。データベース間、またはデータベース内外に多くのデータを移行するのは非常に時間がかかる場合があるため、可能な限り避けることをお勧めします。

注: ODBC データ・ソース名は、各 IBM スポス モデラー ストリームに効果的に埋め込まれます。あるホスト上で作成されたストリームが別のホスト上で実行された場合、データ・ソースの名前はそれぞれのホストで同じである必要があります。また、各入力ノードまたはモデル作成ノードで、「サーバー」タブから異なるデータ・ソースを選択できます。

モデル・オプション

「モデル・オプション」タブで、モデルの名前を指定するか、自動的に名前を生成するかを選択できます。スコアリングのオプションのデフォルト値を設定することもできます。

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

名前が使用されている場合は既存モデルを置換。 このチェック ボックスを選択する場合、同じ名前の既存モデルは上書きされます。

スコアリングで使用可能にする。 モデル・ナゲットのダイアログに表示されるスコアリング・オプションのデフォルト値を設定できます。オプションの詳細については、その特定のナゲットの「設定」タブのヘルプトピックを参照してください。

モデルの管理

スボス モデラー によって IBM Netezza モデルまたは IBM Data Warehouse モデルを構築すると、スボス モデラー 内にモデルが作成されると同時に、IBM Data Warehouse データベース内でもモデルの作成や置き換えが行われます。この種類の スボス モデラー モデルは、データベース・サーバーに格納されているデータベース・モデルの内容を参照します。スボス モデラー では、スボス モデラー モデルと Netezza モデルまたは Data Warehouse モデルの両方に、同一のモデル・キー文字列を生成して格納することで、整合性チェックを実行できます。

Netezza または Data Warehouse の各モデルのモデル名は、「データベース・モデルの一覧表示」ダイアログ・ボックスの「モデル情報」列に表示されます。スボス モデラー モデルのモデル名は、スボス モデラー モデルの「サーバー」タブの「モデル・キー」として表示されます (ストリーム内に配置された場合)。

「検査」ボタンを使用すると、スボス モデラー モデルと Netezza モデルまたは Data Warehouse モデルのモデル・キーが一致するかどうかを検査できます。Netezza または Data Warehouse に同じ名前のモデルがないかモデル・キーが一致しない場合は、スボス モデラー モデルの構築後に Netezza モデルまたは Data Warehouse モデルが削除または再構築されています。

データベース・モデルの一覧表示

スボス モデラー には、IBM Data Warehouse に格納されたモデルを一覧表示するダイアログ・ボックスがあり、そこでモデルを削除できます。このダイアログ・ボックスには、IBM「ヘルパー アプリケーション」ダイアログ・ボックスと、IBM Data Warehouse および IBM Netezza Data Mining の関連ノードの構築、参照、および適用の各ダイアログ・ボックスからアクセス可能です。各モデルについて、次の情報が表示されます。

- モデル名 (リストをソートするのに使用されるモデルの名前)
- 所有者名。
- モデルで使用されるアルゴリズム。
- モデルの現在の状態 (例えば「完了」)。
- モデルが作成された日付。

IBM Data WH 回帰ツリー

回帰ツリーは、数値型対象フィールドの値に基づいて同じ種類のサブセットを派生させるために、ケースのサンプルを繰り返し分割するツリーベースのアルゴリズムです。ディビジョン・ツリーと同様に、回帰ツリーはツリーの葉が十分に小さいか、均一なサブセットに対応するサブセットにデータを分解します。分割は、葉の平均値で合理的に予測できるよう、対象の属性の値のばらつきを減少させるために選択されます。

IBM Data WH 回帰ツリーの作成オプション - ツリーの成長

ツリーの成長とツリーの剪定の作成オプションを設定できます。

ツリーの成長には、以下の作成オプションを使用できます。

最大ツリー深さ。 ルート・ノード下でツリーが成長可能な最大レベル数 (サンプルが再帰的に分割される回数)。デフォルトは 62 で、モデリングのための最大ツリー深度です。

注: モデル・ナゲットのビューアーにモデルがテキスト表示される場合、最大 12 のツリー・レベルが表示されます。

分割の基準。 これらのオプションは、ツリーの分割をいつ停止するかを制御します。デフォルト値を使用しない場合は、「カスタマイズ」をクリックして値を変更します。

- **評価指標の分割。** このクラス評価測定では、ツリーの分割に最適な場所が評価されます。

注: 現在、使用可能なオプションは「分散」のみです。

- **分割の改善度の最小変化量。** 新しい分割がツリーで作成される前に不純度を減少する必要がある最小数。ツリー構築の目標は、似かよった出力値を持つサブグループを作成して、それぞれのノード内にお

ける不純度を最小にすることです。ブランチが適切に分割されて不純度が分割基準によって指定された値を下回ると、ブランチは分割されません。

- **分割の最小インスタンス数。** 分割可能な最小レコード数。分割されていないレコードがこの数より少ない場合、これ以上分割は行われません。このフィールドを使用すると、ツリー内に小さいサブグループが作成されないようにすることができます。

統計。 このパラメーターは、モデルに含まれる統計の数を定義します。次のオプションのいずれかを選択します:

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注: このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「なし」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし:** モデルのスコアリングに必要な統計のみが含まれます。

IBM Data WH ツリーの作成オプション - ツリーの剪定

剪定オプションを使用して、回帰ツリーの剪定基準を指定できます。剪定の目的は、新しいデータに対して予期した精度が改善されない成長しすぎたサブグループを削除することによって、オーバーフィットのリスクを軽減することにあります。

プルーニング測定。 剪定の測定により、ツリーから葉を削除した後、モデルの推定精度が許容限度内に維持されます。次のいずれかを選択できます。

- **MSE** 平均二乗誤差-(デフォルト) 適合線がデータ・ポイントにどの程度近いかを測定します。
- **r2.** R-squared - 回帰モデルによって説明される従属変数の変動の比率です。
- **Pearson:** Pearson の相関係数 - 正規分布されている線形従属変数間の関係の強さを測定します。
- **Spearman:** Spearman の相関係数 - ピアソンの相関に従って弱いと思われるが、実際には強いと考えられる非線形な関係を検出します。

剪定するデータ。 新しいデータに対する想定された精度を推定するには、学習データの一部またはすべてを使用できます。また、指定されたテーブルの別の剪定データセットを使用できます。

- **すべての学習データを使用。** このオプション (デフォルト) は、モデルの精度を推定するためにすべての学習データを使用します。
- **剪定に次のパーセンテージ学習データを使用。** 剪定データに指定した割合で、一方は学習用、一方は剪定用です。

ストリームを実行するごとにデータを同じ方法で区分するようランダム シードを指定する場合、「結果を再現」を選択します。「剪定に使用するシード」フィールドで整数を指定するか、または「生成」をクリックすると、擬似無作為の整数を作成します。

- **既存のテーブルのデータを使用。** モデルの精度を推定するために個別の剪定データセットのテーブル名を指定します。学習データを使用するより信頼性が高いと見なされます。ただし、このオプションにより、学習セットから大きなサブセットのデータが削除され、ディシジョン・ツリーの品質が損なわれます。

Netezza 分裂クラスタリング

分裂クラスタリングは、アルゴリズムが指定された停止ポイントに到達するまで、サブクラスターにクラスターを分割するために繰り返し実行されるクラスター分析の手法です。

クラスター形成は、すべての学習インスタンス (レコード) を含む単一のクラスターから始まります。アルゴリズムの最初の繰り返しでは2つのサブクラスターにデータセットを分割し、後続の反復でそれらをさらにサブクラスターに分割します。停止条件は繰り返しの最大数、データセットが分割されているレベルの最大数、およびさらにデータ区分のためのインスタンスの必要最小限数として指定されます。

結果として得られる階層クラスタリング・ツリーを次の例で使用して、ルート・クラスターからそれらを伝播させることによって、インスタンスを分類することができます。

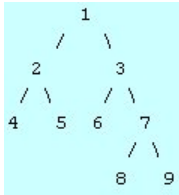


図 3. 分裂クラスタリング・ツリーの例

各レベルで、最良のマッチング・サブクラスターがサブクラスター中心からのインスタンスの距離を基準に選択されます。

葉は負の数に指定されているため、インスタンスが -1（デフォルト）の適用階層レベルでスコアリングされている場合、スコアリングは葉クラスターのみを返します。この例では、クラスター 4、5、6、8、または 9 のいずれかになります。ただし、例えば階層レベルが 2 に設定されている場合、スコアリングはルート・クラスターの下 の 2 番目のレベル (4、5、6、または 7) にあるクラスターの 1 つを返します。

Netezza 分裂クラスタリングのフィールド・オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

「すべて」 ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定の尺度のボタンをクリックして、その測定の尺度のすべてのフィールドを選択します。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

Netezza 分裂クラスタリングの作成オプション

「作成オプション」タブで、モデルを構築するすべてのオプションを設定します。もちろん、「実行」ボタンをクリックすると、すべてデフォルト・オプションのモデルが構築されますが、通常は、それぞれの目的で作成をカスタマイズする必要があります。

距離尺度: 距離が大きいほど、相違が大きくなります。以下のオプションがあります。

- **ユークリッド.** (デフォルト) 2 つのポイント間の距離は、直線で結合することによって計算されます。
- **マンハッタン.** 2 つのポイント間の距離は、それらの座標間の絶対差の合計として計算されます。
- **キャンベラ.** マンハッタン距離と似ていますが、起点に近いデータ・ポイントにより敏感です。
- **最大.** 2 つの点の間の距離は、座標ディメンションに沿った差の最大値として計算されます。

最大反復数: このアルゴリズムは、同じプロセスの複数の反復を実行することによって動作します。指定された回数だけ反復した後、モデルの学習を中止します。

クラスター・ツリーの最大の深さ。 データ・セットを細分化できるレベルの最大数。

結果を再現: 分析を複製できるようになります。整数を指定、または「生成」をクリックすると、擬似無作為の整数を作成します。

分割の最小インスタンス数: 分割可能な最小レコード数。分割されていないレコードがこの数より少ない場合、これ以上分割は行われません。このフィールドを使用して、クラスター・ツリー内の非常に小さいサブグループが作成されないようにすることができます。

IBM Data WH 一般化線型

線型は、数値型入力フィールドの値に基づいてレコードを分類する長きにわたって確立された統計手法です。線型回帰は、予測された出力値と実際の出力値の差異を最小限にする直線または面に適合します。線形モデルは、学習とモデル・アプリケーションにおいて、その単純さにより、さまざまな実世界の現象をモデル化するのに有用である。しかし、線形モデルは、従属(ターゲット)変数の正規分布及び従属変数に対する独立(予測)変数の線形衝突を想定します。

線型が役立つ状況は数多くありますが、上記の想定は適用されません。例えば、様々な製品間における消費者の選択をモデリングする場合、従属変数は多項分布である可能性が高くなります。同様に、年齢に対する収入をモデリングすると、通常は年齢が上がるごとに収入は高くなりますが、2つの間のリンクは直線ほど単純でなくなります。

こうした状況の場合、一般化線形モデルを使用できます。一般化線形モデルは、指定したリンク関数によって従属変数が予測変数に関連するよう、線形モデルを拡張したものです。さらにこのモデルでは、非正規分布にポアソンなどの従属変数を使用できます。

アルゴリズムは、指定された反復数を最大値として、最適なモデルを反復して探します。適合度を計算する場合、従属変数の予測値及び実際の値との間の差の平方和によって誤差が示されます。

IBM Data WH 一般化線形モデルのフィールド・オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 このオプションでは、上流のデータ型ノードまたは上流の入力ノードの「データ型」タブの役割設定(対象や予測など)を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てする場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

「すべて」ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定の尺度のボタンをクリックして、その測定の尺度のすべてのフィールドを選択します。

target. 1つのフィールドを予測の対象として選択します。

レコード ID: 一意のレコード ID として使用されるフィールド。このフィールドの値は各レコードで一意である必要があります(例: カスタマー ID 番号)。

インスタンスの重み。 インスタンスの重みを使用するフィールドを指定します。インスタンスの重みは、入力データの行あたりの重みです。デフォルトでは、すべての入力レコードの重要度は相対的に等しいと想定されています。個々の重みを入力レコードに割り当てることにより、重要度を変更できます。指定するフィールドには、入力データの各行の数値の重みが含まれている必要があります。

予測値(入力): 1つ以上の入力フィールドを選択します。この操作は、データ型ノードのフィールドの役割を「入力」に設定するのと似ています。

IBM Data WH 一般化線形モデル・オプション - 全般

「モデル・オプション」タブで、モデルの名前を指定するか、自動的に名前を生成するかを選択できます。モデル、リンク関数、入力フィールドの反復(あれば)に関連する様々な設定を行い、スコアリング・オプションのデフォルト値を設定します。

モデル名 ターゲットまたはID フィールド(その指定がない場合はモデル タイプ)に基づいてモデル名を生成、またはカスタム名を指定することができます。

フィールド・オプション。 モデル構築用の入力フィールドの役割を指定できます。

一般設定。 これらの設定は、アルゴリズムの停止基準に関連しています。

- **反復の最大数。** 最小値は1、デフォルトは20です。

- **最大エラー (1e)**。アルゴリズムが最適なモデルの検出を停止する必要がある最大誤差値 (指数表記)。最小値は 0、デフォルトは -3、つまり、1E-3 または 0.001 となります。
- **重要でないエラー値のしきい値 (1e)**。エラーがゼロの値を持つものとして扱われる値 (指数表記)。最小値は -1、デフォルトは -7、つまり 1E-7 (または 0.0000001) を下回る誤差の値が有意でないとカウントされます。

分布設定: これらの設定は、従属 (ターゲット) 変数の分布に関連します。

- **応答変数の分布:** 分布のタイプは、**ベルヌーイ** (デフォルト)、**ガウス**、**ポアソン**、**二項分布**、**負の二項分布**、**ワルド** (逆ガウス)、および **ガンマ** のいずれかです。
- **パラメーター**。 (ポアソンおよび二項分布のみ) 「**パラメータの指定**」フィールドに以下のオプションのいずれかを指定する必要があります。

- データからパラメーターを自動的に推定させるには、「**デフォルト**」を選択します。
- 分布疑似尤度の最適化を許可するには、「**疑似(Quasi)**」を選択します。
- パラメーター値を明示的に指定するには、「**明示 (Explicit)**」を指定します。

(二項分布のみ) 二項分布の要求に応じて試行フィールドとして使用される入力テーブル列を指定する必要があります。この列には、二項分布の繰り返し試行数が入ります。

(負の二項分布のみ) デフォルトの -1 を使用するか、または別のパラメーター値を指定することができます。

リンク機能の設定。 これらの設定は、従属変数を予測変数に関連付けるリンク関数に関連しています。

- **リンク関数。** 使用する関数は、**Identity**、**Inverse**、**Invnegative**、**Invsquare**、**Sqrt**、**Power**、**Oddspower**、**Log**、**Clog**、**Loglog**、**Cloglog**、**Logit** (デフォルト)、**Probit**、**Gaussit**、**Cauchit**、**Canbinom**、**Cangeom**、**Cannegbinom** のいずれかです。
- **パラメーター**。 (Power または Oddspower リンク関数のみ) リンク関数が **Power** または **Oddspower** の場合、パラメーター値を指定できます。値を指定するか、デフォルトの 1 を使用するかを選択します。

IBM Data WH 一般化線形モデル・オプション - 相互作用

相互作用パネルには、相互作用を指定するオプション (入力フィールド間の倍数の影響) が表示されます

列の相互作用。 入力フィールド間の交互作用を指定するには、このチェック・ボックスを選択します。相互作用がない場合は、ボックスをオフにします。

ソース・リスト内の 1 つ以上のフィールドを選択し、相互作用リストにドラッグして、相互作用をモデルに入力します。作成する相互作用の種類は、選択項目をドロップするホットスポットによって異なります。

- **MAIN.** ドロップされたフィールドは、交互作用リストの一下部に別の主交互作用として表示されます。
- **2 要因:** ドロップされたフィールドのすべての可能なペアが、2 要因の交互作用として交互作用リストの下部に表示されます。
- **3 要因:** ドロップされたフィールドのすべての可能なトリプレットが、3 要因の交互作用として交互作用リストの下部に表示されます。
- ***** ドロップされたすべてのフィールドの組み合わせは、交互作用リストの下部に単一の交互作用として表示されます。

定数項を含める。 通常、切片はモデルに含まれます。データが原点を通ると仮定できる場合は、切片を除外できます。

ダイアログ・ボックスのボタン

表示の右にあるボタンを使用すると、モデルで使用する用語に変更を加えることができます。



図 4. 「削除」ボタン

削除する項目を選択して削除ボタンをクリックすることで、モデルから項目を削除できます。



図 5. 「並べ替え」 ボタン

並べ替える項目を選択して上下矢印をクリックすることで、モデル内の項目を並べ替えることができます。



図 6. 「カスタム交互作用」 ボタン

カスタム項の追加

$n1 \times x1 \times x1 \times x1$ の形式でカスタム対話を指定できます。フィールド リストからフィールドを選択し、右矢印 ボタンをクリックしてフィールドを **カスタム項目** に追加し、**基準 *** をクリックして次のフィールドを選択し、右矢印 ボタンをクリックします。カスタム相互作用を構築したら、「項目を追加」 をクリックして「相互作用」 パネルに戻ります。

IBM Data WH 一般化線形モデル・オプション - スコアリング・オプション

スコアリングで使用可能にする。モデル・ナゲットのダイアログに表示されるスコアリング・オプションのデフォルト値を設定できます。詳しくは、トピック 78 ページの『IBM Data WH 一般化線形モデル・ナゲット - 「設定」 タブ』を参照してください。

- **入力フィールドを含めます。** モデル出力と予測に入力フィールドを表示する場合は、このチェック・ボックスを選択します。

IBM Data WH ディジション・ツリー

ディジション・ツリーは、分類モデルを示す階層構造です。ディジション・ツリー・モデルによって、学習データのセットから将来の観測値を予測または分類する分類システムを開発できます。分類は、ブランチが分類の分割点を示すツリー構造の形式です。分割は、停止点に達するまでデータをサブグループに繰り返し分割します。停止点のツリーノードは、**葉**と呼ばれます。各葉は、**クラス・ラベル**というラベルを、サブグループのメンバー、またはクラスに割り当てます。

インスタンスの重みとクラスの重み

デフォルトでは、すべての入力レコードとクラスの重要度は相対的に等しいと想定されています。各重みをこれらの項目のいずれかまたは両方のメンバーに割り当てることによって変更できます。例えば、学習データ内のデータ・ポイントがカテゴリ間で現実的に分布していない場合に役に立ちます。重みを使用すると、モデルを偏らせて、データにうまく表れていないカテゴリの補正を行うことができます。対象値の重みが大きくなると、カテゴリの適切な予測の割合が大きくなります。

ディジション・ツリー・モデル作成ノードで、2つの種類の重みを指定できます。**インスタンスの重み**は、入力データの各行に重みを割り当てます。通常、重みはほとんどの場合に 1.0 に指定され、次の表に示すように、過半数よりも重要度が高いか低いケースにのみ高い値または低い値が割り当てられます。

表 5. インスタンスの重み例		
レコード ID	対象	インスタンスの重み
1	drugA	1.1
2	drugB	1.0
3	drugA	1.0
4	drugB	0.3

クラスの重みは、次の表に示すように対象フィールドの各カテゴリーに重みを割り当てます。

表 6. クラスの重み例	
クラス	クラスの重み
drugA	1.0
drugB	1.5

2つのタイプの重みを同時に使用できます。その場合、お互いを掛け合わせ、インスタンスの重みとして使用されます。そのため、2つの前述の例を同時に使用すると、アルゴリズムは次の表に示すようにインスタンスの重みを使用します。

表 7. インスタンスの重みの計算例		
レコード ID	計算	インスタンスの重み
1	1.1×1.0	1.1
2	1.0×1.5	1.5
3	1.0×1.0	1.0
4	0.3×1.5	0.45

Netezza ディジジョン・ツリーのフィールド・オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用: このオプションを選択すると、上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) が使用されます。

カスタム・フィールド割り当てを使用します。 対象、予測、およびその他の役割を手動で割り当てるには、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

リスト内のすべてのフィールドを選択する場合は「すべて」ボタンをクリックし、特定の尺度のすべてのフィールドを選択する場合は各尺度のボタンをクリックします。

target. 1つのフィールドを予測の対象として選択します。

レコード ID: 一意のレコード ID として使用されるフィールド。このフィールドの値は各レコードで一意である必要があります (例: カスタマー ID 番号)。

インスタンスの重み。 ここでフィールドを指定すると、デフォルトのクラスの重み (対象フィールドのカテゴリーごとの重み) の代わりに、またはデフォルトのクラスの重みに加えて、インスタンスの重み (入力データの行ごとの重み) を使用できます。ここで指定するフィールドは、入力データの各行の数値の重みを含むフィールドでなければなりません。詳しくは、トピック [56 ページの『インスタンスの重みとクラスの重み』](#) を参照してください。

予測値 (入力): 1つ以上の入力フィールドを選択します。これは、データ型ノードのフィールドの役割を「入力」に設定する場合と類似しています。

IBM Data WH ディジジョン・ツリーの作成オプション

ツリーの成長には、以下の作成オプションを使用できます。

成長の測定。 これらのオプションは、ツリーの成長を測定する方法を制御します。

- **不純度の測定:** この測定では、ツリーを分割するのに最適な場所が評価されます。これは、サブグループまたはデータのセグメントにおける変動性の測定です。不純度の測定値が低い場合は、グループ内のほとんどのメンバーの基準フィールドまたは対象フィールドの値が類似していることを示します。

サポートされる測定は、「エントロピー」と「Gini」です。これらの測定は、ブランチの所属カテゴリーの確率に基づいています。

- **最大ツリー深さ。** ルート・ノード下でツリーが成長可能な最大レベル数 (サンプルが再帰的に分割される回数)。このプロパティのデフォルト値は 10 であり、このプロパティに設定できる最大値は 62 です。

注: モデル・ナゲットのビューアーにモデルがテキスト表示される場合、最大 12 のツリー・レベルが表示されます。

分割の基準。 これらのオプションは、ツリーの分割をいつ停止するかを制御します。

- **分割の改善度の最小変化量。** 新しい分割がツリーで作成される前に不純度を減少する必要がある最小数。ツリー構築の目標は、似かよった出力値を持つサブグループを作成して、それぞれのノード内における不純度を最小にすることです。ブランチが適切に分割されて不純度が分割基準によって指定された値を下回ると、ブランチは分割されません。
- **分割の最小インスタンス数。** 分割可能な最小レコード数。分割されていないレコードがこの数より少ない場合、これ以上分割は行われません。このフィールドを使用すると、ツリー内に小さいサブグループが作成されないようにすることができます。

統計。 このパラメーターは、モデルに含まれる統計の数を定義します。次のオプションのいずれかを選択します:

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注: このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「なし」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし:** モデルのスコアリングに必要な統計のみが含まれます。

IBM Data WH ディシジョン・ツリー・ノード - クラスの重み

ここでは、重みを各クラスに割り当てることができます。デフォルトでは、1 の値をすべてのクラスに割り当て、重みを等しくします。異なるクラス・ラベルに異なる数値の重みを指定することによって、アルゴリズムが特定のクラスの学習セットに重みを付けるよう指定します。

重みを変更するには、「**重み**」列で変更する重みをダブルクリックし、必要に応じて変更します。

値。 対象フィールドの可能な値から派生した、クラス ラベルのセット。

重量。 特定のクラスに割り当てられる重み。大きな重みをクラスに割り当てると、モデルは、他のクラスと比べてそのクラスに対して敏感になります。

インスタンスの重みと組み合わせてクラスも重みを使用できます。詳しくは、トピック [56 ページの『インスタンスの重みとクラスの重み』](#) を参照してください。

IBM Data WH ディシジョン・ツリー・ノード - ツリーの剪定

剪定オプションを使用して、ディシジョン・ツリーの剪定基準を指定できます。剪定の目的は、新しいデータに対して予期した精度が改善されない成長しすぎたサブグループを削除することによって、オーバーフィットのリスクを軽減することにあります。

剪定の測定。 デフォルトの剪定測定 **正解率** では、ツリーから葉を削除した後、モデルの推定精度が許容限度内に維持されます。剪定を適用しながらクラスの重みを考慮に入れる場合、「**重みつき精度**」を使用します。

剪定するデータ。 新しいデータに対する想定された精度を推定するには、学習データの一部またはすべてを使用できます。また、指定されたテーブルの別の剪定データセットを使用できます。

- **すべての学習データを使用。** このオプション (デフォルト) は、モデルの精度を推定するためにすべての学習データを使用します。
- **剪定に次のパーセンテージ学習データを使用。** 剪定データに指定した割合で、一方は学習用、一方は剪定用です。

ストリームを実行するごとにデータを同じ方法で区分するようランダム シードを指定する場合、「結果を再現」を選択します。「剪定に使用するシード」フィールドで整数を指定するか、または「生成」をクリックすると、擬似無作為の整数を作成します。

- **既存のテーブルのデータを使用。** モデルの精度を推定するために個別の剪定データセットのテーブル名を指定します。学習データを使用するより信頼性が高いと見なされます。ただし、このオプションにより、学習セットから大きなサブセットのデータが削除され、ディシジョン・ツリーの品質が損なわれます。

IBM Data WH 線型

線型モデルは、対象と 1 つ以上の予測値の間の線型関係に基づいて、連続型対象を予測します。直接の線型関係のモデリングのみに限定しながら、線型モデルは比較的単純であり、簡単に解釈されるスコアリング数式を与えます。その適用はより洗練された回帰アルゴリズムで生成されるものに比べて制限されますが、線形モデルは、高速、高効率で使いやすいです。

IBM Data WH 線型作成オプション

「作成オプション」タブで、モデルを構築するすべてのオプションを設定します。もちろん、「実行」ボタンをクリックすると、すべてデフォルト・オプションのモデルが構築されますが、通常は、それぞれの目的で作成をカスタマイズする必要があります。

式を求解するには、特異値分解を使用します。 元の行列の代わりに特異値分解行列を使用すると、数値誤差に対してより頑健であるという利点があり、計算を高速化することもできます。

モデルに切片を含めます。 切片を含めると、解の全体的な精度が向上します。

モデル診断を計算します。 このオプションにより、モデルに対していくつかの診断が計算されます。結果は、後で確認できるようにマトリックスまたは表に保管されます。この診断には、 r^2 乗、残差平方和、分散の推定、標準偏差、 P 値、および T 値が含まれます。

これらの診断は、モデルの有効性と有用性に関連しています。直線性の仮定を満たしていることを保証するために、基礎となるデータに別々の診断を実行する必要があります。

IBM Data WH KNN

最近傍分析は、他のケースに対する類似性に基づいてケースを分類する方法です。この方法は、機械学習において、保存されたパターンやケースへの完全一致を要求することなくデータのパターンを認識する方法として開発されました。類似したケースはお互いに近く、類似していないケースはお互いに離れています。つまり、2 つのケース間の距離は、それらの非類似度の尺度です。

お互いに近いケースは、「近傍」と呼ばれます。新しいケース (ホールドアウト) が表示されたときに、モデル内の各ケースからの距離が計算されます。最も類似した分類 - 最近傍 - が集計され、新しいケースが、最大数の最近傍を含むカテゴリに振り分けられます。

調査する最近隣の数指定でき、この値は k と呼ばれます。写真は、 k の 2 つの異なる値を使用して新しいケースがどのように分類されるかを示しています。 $k=5$ の場合、最近傍の大部分がカテゴリ 1 に属するため、新しいケースはカテゴリ 1 にあります。ただし、 $k=9$ の場合、最近傍の大部分がカテゴリ 0 に属するため、新しいケースはカテゴリ 0 にあります。

また、最近傍分析を使用して、連続型対象の値を計算することもできます。この場合、最近傍の平均対象値または中央対象値を使用して、新しいケースの予測値が取得されます。

IBM Data WH KNN モデル・オプション - 全般

「モデル・オプション - 全般」タブで、モデルの名前を指定するか、自動的に名前を生成するかを選択できます。また、最近傍の数が計算される方法を制御するオプションを設定し、モデルの強化されたパフォーマンスと精度のオプションを設定できます。

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

近傍

距離測定: 距離が大きいほど、相違が大きくなります。以下のオプションがあります。

- **ユークリッド**. (デフォルト) 2つのポイント間の距離は、直線で結合することによって計算されます。
- **マンハッタン**. 2つのポイント間の距離は、それらの座標間の絶対差の合計として計算されます。
- **キャンベラ**. マンハッタン距離と似ていますが、起点に近いデータ・ポイントにより敏感です。
- **最大**. 2つの点の間の距離は、座標ディメンションに沿った差の最大値として計算されます。

最近傍の数 (k): 特定のケースの最近傍の数。使用する近傍の数が多いほど、より正確なモデルが作成されるとは限らないことに注意してください。

k を選択することにより、オーバーフィット（特に「ノイズの多い」データの場合は重要な場合があります）の防止と解決（同様のインスタンスに対する異なる予測）のバランスを制御します。各データセットごとに k の値（値は 1～数十の範囲）を調整する必要があります

パフォーマンスと精度を向上

距離を計算する前に測定を標準化します。 このオプションを選択すると、距離値を計算する前に連続型入力フィールドの測定値が標準化されます。

コア・セットを使用して、大規模なデータ・セットのパフォーマンスを向上させます。 このオプションを選択すると、大規模データ・セットが関係する場合の計算を高速化するために、コア・セット・サンプリングが使用されます。

IBM Data WH KNN モデル・オプション - スコアリング・オプション

「モデル・オプション - スコアリング・オプション」タブでは、スコアリング・オプションのデフォルト値を設定し、個々のクラスに相対的な重みを割り当てることができます。

スコアリングで使用可能にする

入力フィールドを含める: 入力フィールドがデフォルトでスコアリングに含まれるかどうかを指定します。

クラスの重み

モデルを構築する個々のクラスの相対的な重要度を変更したい場合は、このオプションを使用します。

注: このオプションは、分類に KNN を使用している場合にのみ有効です。回帰を実行している場合（つまり、対象フィールドの型が連続型である場合）、オプションは無効になります。

デフォルトでは、1 の値をすべてのクラスに割り当て、重みを等しくします。異なるクラス・ラベルに異なる数値の重みを指定することによって、アルゴリズムが特定のクラスの学習セットに重みを付けるよう指定します。

重みを変更するには、「**重み**」列で変更する重みをダブルクリックし、必要に応じて変更します。

値. 対象フィールドの可能な値から派生した、クラス ラベルのセット。

重量. 特定のクラスに割り当てられる重み。大きな重みをクラスに割り当てると、モデルは、他のクラスと比べてそのクラスに対して敏感になります。

IBM Data WH K-Means

K-Means ノードは、クラスター分析の方法を提供する k -means アルゴリズムを実行します。このノードを使用して、データ・セットをグループにクラスター化できます。

アルゴリズムは距離に基づくアルゴリズムで、距離メトリック（関数）に依存してデータ・ポイント間の類似性を測定します。データ・ポイントは、使用する距離メトリックに従って最も近いクラスターに割り当てられます。

アルゴリズムは、同じ基本プロセスを何度か反復することによって実行します。学習インスタンスは最も近いクラスターに割り当てられます（指定された距離関数に関しては、インスタンスとクラスター中心に適用されます）。すべてのクラスター中心は、特定のクラスターに割り当てられたインスタンスの平均属性値のベクトルとして再計算されます。

IBM Data WH K-Means のフィールド・オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

「すべて」ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定の尺度のボタンをクリックして、その測定の尺度のすべてのフィールドを選択します。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

IBM Data WH K-Means の作成オプション・タブ

作成オプションを設定することにより、用途に合わせてモデルの作成をカスタマイズできます。

デフォルト・オプションを使用してモデルを作成する場合は、「実行」をクリックします。

距離尺度: このパラメーターは、データ・ポイント間の距離の測定方法を定義します。距離が大きくなると、非類似度も大きくなります。次のオプションのいずれかを選択します:

- **ユークリッド。** ユークリッド測定は、2 つのデータ・ポイント間の直線距離です。
- **正規化ユークリッド (Normalized Euclidean)。** 正規化ユークリッド測定はユークリッド測定に類似していますが、平方標準偏差によって正規化されます。ユークリッド測定とは異なり、正規化ユークリッド測定はスケール不変でもあります。
- **Mahalanobis。** Mahalanobis 測定は、入力データの相関を考慮に入れた一般化ユークリッド測定です。正規化ユークリッド測定と同様に、Mahalanobis 測定はスケール不変です。
- **Manhattan。** Manhattan 測定は 2 つのデータ・ポイント間の距離で、それらの座標間の絶対距離の合計として計算されます。
- **Canberra。** Canberra 測定は Manhattan 測定に似ていますが、原点に近いデータ・ポイントに対してより感度が高くなります。
- **最大 最大測定は、** 2 つのデータ・ポイント間の距離で、座標の次元に沿ったそれらの相違の最大値として計算されます。

クラスター数: このパラメーターは、作成するクラスターの数进行を定義します。

反復の最大数。 アルゴリズムは、同じプロセスを何度か反復します。このパラメーターは、モデルの学習を停止する前の反復数を定義します。

統計。 このパラメーターは、モデルに含まれる統計の数を定義します。次のオプションのいずれかを選択します:

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注: このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「なし」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし:** モデルのスコアリングに必要な統計のみが含まれます。

結果を再現。 ランダム・シードを設定して分析を複製する場合は、このチェック・ボックスを選択します。整数を指定するか、「生成」をクリックして整数の疑似乱数を作成します。

IBM Data WH Naive Bayes

Naive Bayes は、分類の問題に対応する有名なアルゴリズムです。提示されたすべての予測変数は相互に依存関係がないものとして処理されるので、モデルは *naïve* と命名されます。Naive Bayes は拡張性のある高速のアルゴリズムであり、複数の属性と対象属性の組み合わせに対して、条件に関連する確率を計算します。学習データから、個別の確率が計算されます。各入力変数の各値カテゴリーを計算単位とすると、この確率は各対象クラスの確率を表します。

Netezza ベイズ・ネットワーク

Bayesian Network (ベイズ) は、モデルで、データセットの変数およびそれらの間の確率的、または条件付き独立性が表示されます。Netezza ベイズ・ノードを使用すると、観測された情報および記録された情報を常識という実際の知識を組み合わせることによって確率モデルを作成し、表面的にはリンクしていない属性を使用して発生の尤度を確立できます。

Netezza Bayes ネットワークのフィールド・オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

このノードの場合、対象フィールドはスコアリングのためにのみ必要ですので、このタブには表示されません。このノードの「モデルのオプション」タブで、またはモデルナゲットの「設定」タブで、データ型ノード上で対象を設定または変更することができます。詳しくは、トピック [72 ページの『Netezza Bayes Net ナゲット - 「設定」タブ』](#) を参照してください。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

「すべて」 ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定の尺度のボタンをクリックして、その測定の尺度のすべてのフィールドを選択します。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

Netezza Bayes ネットワークの作成オプション

「作成オプション」タブで、モデルを構築するすべてのオプションを設定します。もちろん、「実行」ボタンをクリックすると、すべてデフォルト・オプションのモデルが構築されますが、通常は、それぞれの目的で作成をカスタマイズする必要があります。

基本索引。 内部管理を容易にするために最初の属性 (入力フィールド) に割り当てられる数値 ID。

サンプル・サイズ。 属性の数が多く、非常に長い処理時間を引き起こすことがある場合に採用されるサンプルのサイズ。

実行時に追加の情報を表示: このボックスをオンにすると (デフォルト)、追加情報がメッセージ ダイアログ・ボックスに表示されます。

Netezza 時系列

時系列は、毎日の株価や毎週の営業データなど、連続する (必ずしも正規ではない) 時点で測定した数値データ値のシーケンスです。傾向や季節性 (繰り返しパターン) などの行動を強調したり、過去の事象から将来の行動を予測する場合などに、このようなデータの分析が役立ちます。

Netezza 時系列は、次の時系列アルゴリズムをサポートします。

- スペクトル解析

- 指数平滑法
- 自己回帰和分移動平均 (ARIMA)
- 季節的傾向分解

これらのアルゴリズムは、時系列を傾向と季節成分に分類します。予測に使用できるモデルを構築するために、これらのコンポーネントを分析します。

スペクトル解析を使用して時系列の定期的な動作を識別します。複数の基底の周期で構成された時系列の場合、または相当数のランダムノイズがデータ内に存在する場合、スペクトル解析により、最も明らかに周期的なコンポーネントを特定することができるようになります。この方法により、周波数領域の系列に時間領域からシリーズを変換することによって、定期的な動作の周波数を検出します。

指数平滑化は、前の時系列の観測結果に重み付けされた値を使用して将来の値を予測する方法です。指数平滑化により、観測値の影響は指数関数的な方法で時間の経過に従って短くなります。このメソッドは、新しいデータが入ってくるごとに予測を調整、その追加、傾向、季節性を考慮し、一度に1つのポイントを予想します。

ARIMA モデルは、指数平滑化モデルより洗練された方法で傾向および季節性のコンポーネントをモデリングします。これは、差異の程度と同様に自己回帰および移動平均の順序を明示して指定することと関連します。

注: 実際面では、郵送するカタログの数または会社の Web ページのヒット数など予測対象の一連の性質を説明する上で役立つ予測値を含める場合は、ARIMA モデルが最も有用です。指数平滑化モデルは、性質や傾向の理由を理解しようとしなくて、時系列の性質や傾向を記述します。

季節的傾向分解では、傾向分析を実行するために、時系列からの定期的な動作を削除し、二次関数などのトレンドの基本的な形状を選択します。これらの基本的な形状は、値が残差の平均二乗誤差（つまり、時系列の近似と実測値の差）を最小化するように決定されるさまざまなパラメーターがあります。

Netezza 時系列の値の補間

補間は、時系列データの欠損値の推定および挿入のプロセスです。

時系列の間隔が定期的で、いくつかの値が存在しない場合、欠損値は線形補間を用いて推定することができます。空港での毎月の乗客数について考えてみましょう。

表 8. 空港ターミナルの月間利用者数	
月	乗客
3	3,500,000
4	3,900,000
5	-
6	3,400,000
7	4,500,000
8	3,900,000
9	5,800,000
10	6,000,000

この場合、線形補間は月 5 の欠損値を 3,650,000 (月 4 と月 6 の中間点) と推定します。

不定期な間隔は別の方法で処理されます。次に、温度読み取りについて考えてみましょう。

表 9. 温度読み取り		
日付	時間	温度
2011-07-24	7:00	57

表 9. 温度読み取り (続き)		
日付	時間	温度
2011-07-24	14:00	75
2011-07-24	21:00	72
2011-07-25	7:15	59
2011-07-25	14:00	77
2011-07-25	20:55	74
2011-07-27	7:00	60
2011-07-27	14:00	78
2011-07-27	22:00	74

ここでは、3 日間で 3 つのポイントの読み取り値がありますが、さまざまな場合で、そのうちのいくつかだけが数日の間で共通です。また、数日のうち 2 日だけが連続します。

この状況は、集計の計算、ステップ・サイズの決定のいずれかの方法で処理できます。

データ集計は、データの意味的情報に基づいた式に従って計算された毎日の集計です。これにより、次のようなデータセットが生成されます。

表 10. 温度読み取り (集計)		
日付	時間	温度
2011-07-24	24:00	69
2011-07-25	24:00	71
2011-07-26	24:00	NULL
2011-07-27	24:00	72

また、このアルゴリズムは、異なる時系列を異なるものとして扱い、適切なステップ・サイズを決定することができます。この場合、アルゴリズムによって決定されたステップ・サイズは 8 時間となり、次のようになります。

表 11. ステップ・サイズを計算した温度の読み取り		
日付	時間	温度
2011-07-24	6:00	
2011-07-24	14:00	75
2011-07-24	22:00	
2011-07-25	6:00	
2011-07-25	14:00	77
2011-07-25	22:00	
2011-07-26	6:00	
2011-07-26	14:00	
2011-07-26	22:00	
2011-07-27	6:00	
2011-07-27	14:00	78

表 11. ステップ・サイズを計算した温度の読み取り (続き)		
日付	時間	温度
2011-07-27	22:00	74

ここでは、4つの測定値だけが、元の測定に対応していますが、元の時系列の他の値によって、欠損値が再度補間により計算できるようになります。

Netezza 時系列フィールド・オプション

「フィールド」タブで、ソース・データの入力フィールドの役割を指定します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

target. 1つのフィールドを予測の対象として選択します。尺度が連続型のフィールドでなければなりません。

(予測値) 時点。 (必須) 時系列の日付または時刻の値を含む入力フィールド。このフィールドには、測定の尺度が連続型またはカテゴリー型のフィールド、および日付、時間、タイム・スタンプ、または数値のデータ・ストレージ・タイプを指定する必要があります。ここで指定したフィールドのデータ・ストレージ・タイプは、このモデル作成ノードの他のタブでいくつかのフィールドの入力タイプも定義します。

(予測値) 時系列 ID (単位)。 入力に複数の時系列が含まれる場合に使用します。

Netezza 時系列構築オプション

構築オプションには2つのレベルがあります。

- 基本 - アルゴリズムの選択、補間、使用する時間の範囲の設定
- 高度 - 予測の設定

このセクションでは、基本オプションについて説明します。

「作成オプション」タブで、モデルを構築するすべてのオプションを設定します。もちろん、「実行」ボタンをクリックすると、すべてデフォルト・オプションのモデルが構築されますが、通常は、それぞれの目的で作成をカスタマイズする必要があります。

アルゴリズム(A)

これらは、使用する時系列アルゴリズムに関連する設定です。

アルゴリズム名。 使用する時系列アルゴリズムを選択します。利用可能なアルゴリズムは、「**スペクトル解析**」、「**指数平滑化**」(デフォルト)、「**ARIMA**」、または「**季節的傾向分解**」です。詳しくは、トピック 62 ページの『[Netezza 時系列](#)』を参照してください。

傾向。 (指数平滑法のみ) 時系列がトレンドを示す場合、単純な指数平滑法はうまく実行されません。このフィールドを使用して傾向があれば指定し、アルゴリズムがその傾向を考慮できるようにします。

- **システム決定。** (デフォルト) システムは、このパラメーターの最適な値を見つけようとします。
- **なし (N)。** 時系列は傾向を示しません。
- **加算 (A)。** 時間の経過とともに着実に増加する傾向。
- **減衰付加 (DA)。** 最終的に消滅する加法トレンド。
- **乗法 (M)。** 時間の経過とともに増加する傾向。通常は、安定した加法的傾向よりも急速に増加します。
- **Damped Multiplicative (DM)。** 最終的に消滅する乗法的傾向。

季節性。 (指数平滑法のみ) このフィールドを使用して、時系列がデータ内の季節性パターンを示すかどうかを指定します。

- **システム決定。** (デフォルト) システムは、このパラメーターの最適な値を見つけようとします。
- **なし (N)。** 時系列では季節性のパターンは表示されません。

- **加算 (A)**。季節変動のパターンは、時間の経過とともに安定した上昇傾向を示しています。
- **乗法 (M)**。加法的季節性と同じですが、さらに季節変動の振幅 (高い点と低い点の間の距離) が、変動の全体的な上昇傾向に対して相対的に増加します。

ARIMA のシステム決定設定を使用します。 (ARIMA のみ) ARIMA アルゴリズムの設定をシステムに決定させる場合は、このオプションを選択します。

指定します。 (ARIMA のみ) このオプションを選択し、ボタンをクリックして ARIMA 設定を手動で指定します。

補間

時系列ソース・データに欠損値がある場合、推定値を挿入する方法を選択して、データの隙間を埋めます。詳しくは、トピック [63 ページ](#)の『[Netezza 時系列の値の補間](#)』を参照してください。

- **線形**。時系列の間隔が規則的であるが、いくつかの値が存在しない場合、この方法を選択します。
- **指数スプライン**。既知のデータ・ポイント値が高い率で増減する平滑化曲線に適合します。
- **立方スプライン**。平滑化曲線を既知のデータ・ポイントに適合させ、欠損値を推定します。

時間範囲

ここでは、時系列の完全なデータの範囲、またはそのデータの連続したサブセットを使用してモデルを作成するかどうかを選択できます。これらのフィールドの有効な入力は、「フィールド」タブの「時点」で指定されたフィールドのデータ・ストレージ・タイプによって定義されます。詳しくは、トピック [65 ページ](#)の『[Netezza 時系列フィールド・オプション](#)』を参照してください。

- **データで使用可能な最も早い時刻と最も遅い時刻を使用します。** 時系列データの全範囲を使用する場合は、このオプションを選択します。
- **時間枠を指定します。** 時系列の一部のみを使用する場合は、このオプションを選択します。「**最も早い時間 (開始)**」と「**最も遅い時間 (終了)**」フィールドを使用して境界を指定します。

ARIMA 構造

ARIMA モデルの様々な非季節性および季節性要素の値を指定します。いずれの場合も、演算子を = (等しい) または <= (以下) に設定してから、隣接するフィールドに値を指定します。すべての値は、程度を示す負でない整数にする必要があります。

非季節性。 モデルのさまざまな非季節性の成分の値。

- **自己相関度 (p)**。モデル内の自己回帰の次数の数値です。自己回帰の次数は、系列のどの過去の値を使用して現在の値を予測するかを指定します。例えば、自己回帰の次数 2 は、現在の値を予測するために系列の値を過去の 2 期間使用するように指定します。
- **誘導 (d)**。モデルを推定する前に系列に適用される差分の次数を指定します。差分は、トレンドが存在する場合に必要となり (通常、トレンドの存在する系列は非定常性であり、ARIMA モデル作成は定常性を前提とするため)、トレンドの影響を除去する目的に使用されます。差分の次数は、系列のトレンドの次数に対応しています (1 次差分は線型トレンドを表し、2 次差分は 2 次トレンドを表す、など)。
- **移動平均 (q)**。モデル内の移動平均の次数の数値。移動平均の次数は、過去の値の系列平均からの偏差を、現在の値の予測にどのように使用するかを指定します。例えば、移動平均の次数 1 および 2 は、系列の現在の値を予測するときに、最近の 2 期間のそれぞれから取得した系列の平均値の偏差が考慮されるように指定します。

季節性。 季節性自己相関 (SP)、季節性導出 (SD)、季節性移動平均 (SQ) の各要素は、対応する非季節性の要素と同じ役割を果たします。ただし、季節次数の場合は、現在の系列値が 1 つ以上の季節期間で区切られた過去の系列値の影響を受けます。例えば、月単位のデータ (季節期間 12) の場合、季節次数 1 は、現在の期間より 12 期間前の系列値の影響を受けることを意味します。月単位のデータに関しては、季節次数 1 は、非季節次数 12 を指定するのと同様です。

季節性の設定は、季節性がデータで検出された場合、または「詳細」タブの「期間」設定を指定した場合にのみ考慮されます。

Netezza 時系列構築オプション - 詳細設定

詳細設定を使用して、予測のオプションを指定できます。

モデル作成オプションのシステム決定設定を使用します。システムで拡張設定を決定する場合は、このオプションを選択します。

指定します。 拡張オプションを手動で指定する場合は、このオプションを選択します。(アルゴリズムがスペクトル解析の場合、このオプションは無効です)

- **期間/期間の単位。** その時間が経過すると、時系列のいくつかの特性的な動作が繰り返される期間。例えば、週間営業成績の時系列について、期間に 1 を指定し、単位に「週」を指定します。「期間」には、負ではない整数を指定する必要があります。「期間の単位」には、「ミリ秒」、「秒」、「分」、「時間」、「日」、「週」、「四半期」、「年」のいずれかを指定します。「期間」が設定されている場合、時間のタイプが数値でない場合は「期間の単位」を指定しないでください。ただし、「期間」を指定する場合、「期間の単位」を指定する必要があります。

予測の設定。 特定の時点まで、または特定の時点まで、予測を作成することを選択できます。これらのフィールドの有効な入力値は、「フィールド」タブの「時点」で指定されたフィールドのデータ・ストレージ・タイプによって定義されます。詳しくは、トピック 65 ページの『Netezza 時系列フィールド・オプション』を参照してください。

- **予測期間。** 予測のエンドポイントのみを指定する場合は、このオプションを選択します。予測はこの時点まで行われます。
- **予測時間。** 予測を行う 1 つ以上の時点指定するには、このオプションを選択します。「追加」をクリックして、時点のテーブルに新しい行を追加します。行を削除するには、行を選択して「削除」をクリックします。

Netezza 時系列のモデル・オプション

「モデル・オプション」タブで、モデルの名前を指定するか、自動的に名前を生成するかを選択できます。モデル出力のオプションのデフォルト値を設定することもできます。

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

スコアリングで使用可能にする。 モデル・ターゲットのダイアログに表示されるスコアリング・オプションのデフォルト値を設定できます。

- **結果に過去の値を含める:** デフォルトでは、モデル出力に過去のデータ値 (予測に使用された値) は含まれません。これらの値を含めるにはこのチェックボックスをオンにします。
- **結果に補間された値を含めます。** 出力に履歴値を含めることを選択した場合、補間された値があればそれも含めるには、このボックスを選択します。補間は過去のデータにのみ機能するため、「結果に過去の値を含める」が選択されていない場合このボックスは無効です。詳しくは、トピック 63 ページの『Netezza 時系列の値の補間』を参照してください。

IBM Data WH TwoStep

TwoStep ノードは、大規模データセットにわたってデータをクラスタ化する方法を提供する TwoStep アルゴリズムを実装します。

このノードを使用すると、使用可能なリソース (例えば、メモリーや時間の制約) を考慮しながら、データをクラスタ化できます。

TwoStep アルゴリズムは、次の方法でデータをクラスタ化するデータベースマイニングアルゴリズムです。

1. クラスタ機能 (CF) ツリーが作成されます。このバランスに優れたツリーは、類似の入力レコードが同じツリー ノードの一部となる階層クラスタリングのためのクラスタ機能を格納します。
2. CF ツリーの葉は、最終クラスタリングの結果を生成するためにメモリー内で階層的にクラスタ化されます。クラスタの最適な数は自動的に決定されます。クラスタの最大数を指定する場合、指定された制限内のクラスタの最適な数が決定されます。

3. クラスタリングの結果は、K-Means アルゴリズムに似たアルゴリズムがデータに適用される 2 番目のステップで調整されます。

IBM Data WH TwoStep フィールド・オプション

フィールド オプションを設定することにより、上流ノードで定義されているフィールドの役割の設定を指定できます。フィールド割り当てを手動で行うこともできます。

項目の選択: 上流のデータ型ノードまたは上流の入力ノードの「データ型」タブの役割設定を使用するには、このオプションを選択します。役割設定は、例えば対象や予測です。

カスタム・フィールド割り当てを使用します。 対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印を使用して、このリストの項目を右側の役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

IBM Data WH TwoStep 作成オプション

作成オプションを設定することにより、用途に合わせてモデルの作成をカスタマイズできます。

デフォルト・オプションを使用してモデルを作成する場合は、「**実行**」をクリックします。

距離測定: このパラメーターは、データ・ポイント間の距離の測定方法を定義します。距離が大きくなると、非類似度も大きくなります。以下のオプションがあります。

- **対数尤度。** この尤度測定により、変数の確率分布を求めます。連続型変数は正規分布しているものと仮定され、カテゴリ変数は多項分布しているものと仮定されます。すべての変数は独立しているものと仮定されます。
- **ユークリッド。** ユークリッド測定は、2 つのデータ・ポイント間の直線距離です。
- **正規化ユークリッド (Normalized Euclidean)。** 正規化ユークリッド測定はユークリッド測定に類似していますが、平方標準偏差によって正規化されます。ユークリッド測定とは異なり、正規化ユークリッド測定はスケール不変でもあります。

クラスター番号。 このパラメーターは、作成するクラスターの数进行を定義します。以下のオプションがあります。

- **クラスターを自動的に計算します。** クラスター数が自動的に計算されます。「**最大**」フィールドで最大クラスター数を指定できます。
- **クラスター数を指定します。** 作成するクラスター数を指定します。

統計。 このパラメーターは、モデルに含まれる統計の数を定義します。以下のオプションがあります。

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注: このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「**なし**」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし:** モデルのスコアリングに必要な統計のみが含まれます。

結果を再現。 ランダム・シードを設定して分析を複製する場合は、このチェック・ボックスを選択します。整数を指定するか、「**生成**」をクリックして整数の疑似乱数を作成します。

IBM Data WH PCA

主成分分析 (PCA) は、データの複雑さを軽減するために設計された強力なデータ削減技術です。入力フィールドの線型結合が検出されます。成分が互いに直交する (相関しない) 場合に、フィールドのセット全体の分散を把握するのに役立ちます。どちらの手法でも、元の入力フィールド・セットの情報を効果的に要約する少数の派生フィールド (主成分) の検出が目標です。

注：小文字のフィールド名が使用されている場合、モデルのスコアリング時にエラーが発生することがあります。これは、Db2 Data Warehouse の既知の不具合です。回避策としては、スコアリング前にすべてのフィールドを大文字に名前変更してください。

IBM Data WH PCA フィールド・オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定 of 尺度を示します。

「すべて」ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定 of 尺度のボタンをクリックして、その測定 of 尺度のすべてのフィールドを選択します。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

IBM Data WH PCA 作成オプション

「作成オプション」タブで、モデルを構築するすべてのオプションを設定します。もちろん、「実行」ボタンをクリックすると、すべてデフォルト・オプションのモデルが構築されますが、通常は、それぞれの目的で作成をカスタマイズする必要があります。

PCA を計算する前にデータをセンター化します。 チェック・マークが付いている場合 (デフォルト)、このオプションは、分析の前にデータ・センタリング (「平均減算」とも呼ばれます) を実行します。データ・センタリングは、第 1 主成分が最大分散の方向を記述することを保証するために必要です。そうでない場合、成分がデータの平均値により密接に対応している場合があります。データがすでにこの方法で調製されている場合、通常パフォーマンスの向上のためだけにこのオプションを無効にします。

PCA を計算する前にデータ・スケーリングを実行します。 このオプションは、分析の前にデータ・スケーリングを実行します。そうすることで、別の変数が異なる単位で測定されるとき、分析が恣意的でないようにします。最も単純な形式のデータのスケールリングは、その標準偏差で各変数を分割することによって実現できます。

PCA の計算には、正確ではないが高速な方法を使用してください。 このオプションを使用すると、アルゴリズムは、より正確ではあるが高速な、基本コンポーネントの検出方法 (forceEigensolve) を使用します。

IBM Data WH モデルおよび Netezza モデルの管理

IBM Data Warehouse モデルおよび IBM NetezzaAnalytics モデルは、他の IBM スポス モデラー モデルのようにキャンバスおよび「モデル」パレットに追加され、ほぼ同じように使用できます。ただし、IBM スポス モデラー で生成された各 IBM Data Warehouse モデルまたは IBM NetezzaAnalytics モデルは実際にはデータベース・サーバー上にあるモデルへの参照であるなどの、重要な違いがいくつかあります。ストリームを正しく機能させるためには、モデルが作成されたデータベースに接続する必要があります。モデルテーブルを外部プロセスにより変更してはいけません。

IBM Data Warehouse モデルと IBM NetezzaAnalytics モデルのスコアリング

モデルは、キャンバス上で金色のモデル・ナゲット・アイコンで示されます。ナゲットの主な目的は、データをスコアリングし、予測を生成、またはモデルのプロパティの詳細な分析を可能にすることです。スコアは、このセクションで後述するように、ナゲットにテーブル・ノードを接続し、ストリームのそのブランチを実行することによって見えるようにすることができます、1 つ以上の追加データ・フィールドの形式で追加されます。ディビジョンツリーや回帰ツリーのダイアログ・ボックスなど、いくつかのナゲットのダイアログボックスには、さらにモデルの視覚的表現を提供する「モデル」タブがあります。

追加フィールドは、対象フィールドの名前に追加された接頭辞 \$<id>- によって区別されます。<id> はモデルによって異なり、追加される情報の種類を識別します。さまざまな識別子について、それぞれのモデル・ナゲットのトピックで説明されています。

スコアを表示するには、次の手順を実行します。

1. テーブル・ノードをモデル・ナゲットに接続します。
2. テーブル・ノードを開きます。
3. 「実行 (Run)」をクリックします。
4. テーブル出力ウィンドウの右側にスクロールし、追加フィールドとそれらのスコアを表示します。

IBM Data WH と Netezza のモデル・ナゲットの「サーバー」タブ

「サーバー」タブでは、モデルのスコアリングサーバのオプションを設定することができます。上流に指定されたサーバの接続を使い続けるか、ここで指定した別のデータベースにデータを移動することができます。

IBM Data Warehouse サーバーの詳細: ここで、モデルに使用するデータベースの接続の詳細を指定します。

- **上流の接続を使用:** (デフォルト) データベース入力ノードなど、上流のノードで指定した接続の詳細を使用します。このオプションは、すべての上流ノードが SQL プッシュバックを使用できる場合にのみ有効です。この場合、SQL が完全にすべての上流ノードを実装するため、データベース外のデータを移動する必要はありません。
- **接続へのデータの移動:** ここで指定したデータベースにデータを移動します。データを移動することにより、データが別の IBM Data Warehouse データベース、他のベンダーのデータベースにある場合、またはデータがフラット・ファイルにある場合でもモデリングを実行できるようにします。また、ノードが SQL プッシュバックを実行していないためデータが抽出されていない場合、データはここで指定されたデータベースに戻されます。接続を参照して選択するには、「編集」ボタンをクリックします。



注意: 通常、IBM NetezzaAnalytics と IBM Data Warehouse は、非常に大きいデータ・セットで使用されます。データベース間、またはデータベース内外に多くのデータを移行するのは非常に時間がかかる場合があるため、可能な限り避けることをお勧めします。

モデル名。 モデルの名前。この名前は通知専用に表示されます。ここで名前を変更することはできません。

IBM Data WH ディシジョン・ツリー・モデル・ナゲット

ディシジョン・ツリー・モデル・ナゲットでは、モデリング操作の出力を表示し、モデルをスコアリングするためのオプションを設定することもできます。

ディシジョン・ツリー・モデル・ナゲットを含むストリームを実行すると、デフォルトでは、ノードは 1 つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 12. ディシジョン・ツリーのモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$I-target_name	現在のレコードの予測値。

モデル作成ノードまたはモデル・ナゲットで「レコードのスコアリングに割り当てられたクラスの確率を計算する」を選択した場合、さらにフィールドを追加します。

表 13. ディシジョン・ツリーのモデルスコアリング・フィールド - 追加	
追加フィールドの名前	意味
\$IP-target_name	予測の確信度 (0.0 ～ 1.0)。

IBM Data WH ディシジョン・ツリー・ナゲット - 「モデル」 タブ

「モデル」タブには、ディシジョン・ツリー・モデルの「予測値の重要度」がグラフィカル形式で表示されます。棒の長さは、予測値の重要度を表しています。

注：IBM Netezza Analytics バージョン 2.x 以前で作業している場合、ディシジョン・ツリー・モデルの内容はテキスト形式でのみ表示されます。

これらのバージョンの場合、以下の情報が表示されます。

- テキストの各行は、ノードまたは葉に対応しています。
- インデントはツリー・レベルを反映します。
- ノードの場合、分割の条件が表示されます。
- 葉の場合、割り当てられたクラスのラベルが表示されます。

IBM Data WH ディシジョン・ツリー・ナゲット- 「設定」 タブ

「設定」タブで、モデルをスコアリングするためのオプションを設定できます。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

レコードのスコアリングのために割り当てられたクラスの確率を計算する. (ディシジョン・ツリーと Naive Bayes のみ) 選択した場合、このオプションは、追加のモデリング・フィールドに確信 (確率) フィールドと予測フィールドが含まれることを意味します。このチェックボックスをオフにすると、予測のみのフィールドが生成されます。

決定論的入力データを使用します。 このオプションが選択された場合、同じビューの複数のパスを実行するすべての Netezza アルゴリズムが各パスに確実に同じデータセットを使用するようになります。非決定的データが使用されていることを示すためにこのチェックボックスをクリアすると、データ区分ノードによって作成されるものなど、処理するデータ出力を保持する一時テーブルが作成されます。このテーブルは、モデルの作成後に削除されます。

IBM Data WH ディシジョン・ツリー・ナゲット- 「ビューア」 タブ

「ビューア」タブには、SPSS Modeler でディシジョン ツリー モデルが表示されるのと同じ方法で、ツリーモデルのツリー プレゼンテーションが表示されます。

注：モデルが IBM Netezza Analytics バージョン 2.x 以前で構築されている場合、「ビューア」タブは空です。

IBM Data WH K-Means モデル・ナゲット

K-Means モデル・ナゲットには、クラスター化モデルが取得したすべての情報と、学習データと推定プロセスに関する情報が含まれます。

K-Means モデル・ナゲットを含むストリームを実行すると、そのレコードの所属クラスターと割り当てられたクラスターの中心からの距離を含む 2 つの新規フィールドが追加されます。\$KM-K-Means という名前の新規フィールドは所属クラスター用で、\$KMD-K-Means という名前の新規フィールドはクラスターの中心からの距離用です。

IBM Data WH K-Means ナゲット - 「モデル」 タブ

「モデル」タブには、クラスターのフィールドの要約統計量および分布を表示する各種のグラフィック表示があります。モデルからデータをエクスポートしたり、ビューをグラフィックとしてエクスポートしたりすることができます。

IBM Netezza Analytics バージョン 2.x 以前で作業している場合、あるいは距離測度として Mahalanobis を使用してモデルを構築する場合は、K-Means モデルの内容はテキスト形式でのみ表示されます。

これらのバージョンの場合、以下の情報が表示されます。

- **要約統計量。** 最小クラスターおよび最大クラスターの両方について、統計の要約にレコード数が表示されます。「統計の要約」には、これらのクラスターによって占有されるデータ・セットの割合も表示されます。またリストには、最大クラスターから最小クラスターに対するサイズ比が表示されます。
- **クラスタリング要約。** クラスタリング要約には、アルゴリズムで作成されたクラスターがリストされます。各クラスターについて、テーブルにはクラスターのレコード数、これらのレコードのクラスター中心からの平均距離が表示されます。

IBM Data WH K-Means ナゲット - 「設定」 タブ

「設定」タブで、モデルをスコアリングするためのオプションを設定できます。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

距離測定: 距離が大きいほど、相違が大きくなります。以下のオプションがあります。

- **ユークリッド.** (デフォルト) 2つのポイント間の距離は、直線で結合することによって計算されます。
- **マンハッタン.** 2つのポイント間の距離は、それらの座標間の絶対差の合計として計算されます。
- **キャンベラ.** マンハッタン距離と似ていますが、起点に近いデータ・ポイントにより敏感です。
- **最大.** 2つの点の間の距離は、座標ディメンションに沿った差の最大値として計算されます。

Netezza Bayes ネットワークのモデル・ナゲット

Bayes Net のモデルナゲットは、モデルをスコアリングするためのオプションを設定する手段を提供します。

Bayes Net モデル・ナゲットを含むストリームを実行すると、ノードは1つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 14. Bayes Net のモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$BN-target_name	現在のレコードの予測値。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

Netezza Bayes Net ナゲット - 「設定」 タブ

「設定」タブでは、モデルのスコアリングのオプションを設定することができます。

target. 現在の対象とは異なる対象フィールドをスコアリングしたい場合は、ここで新しい対象を選択します。

レコード ID: 「レコード ID」フィールドが指定されていない場合は、ここで使用するフィールドを選択してください。

予測のタイプ。 使用する予測アルゴリズムのバリエーションは、以下のとおりです。

- **最適 (相関度が最も高い隣接)。** (デフォルト) 最も相関性の高い隣接ノードを使用します。
- **近傍 (近傍の重み付き予測)。** すべての隣接ノードの重み付き予測を使用します。
- **NN-隣接 (NULL 以外の隣接)。** NULL 値を持つノード (つまり、予測が計算されるインスタンスの欠損値を持つ属性に対応するノード) を無視する点を除き、前のオプションと同じです。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

IBM Data WH Naive Bayes のモデル・ナゲット

Naive Bayes のモデルナゲットは、モデルをスコアリングするためのオプションを設定する手段を提供します。

Naive Bayes のモデル・ナゲットを含むストリームを実行すると、デフォルトでは、ノードは1つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 15. Naive Bayes のモデルスコアリング・フィールド - デフォルト	
追加フィールドの名前	意味
\$I-target_name	現在のレコードの予測値。

モデル作成ノードまたはモデル・ナゲットで「**レコードのスコアリングに割り当てられたクラスの確率を計算する**」を選択した場合、さらにフィールドを2つ追加します。

表 16. Naive Bayes のモデルスコアリング・フィールド - 追加	
追加フィールドの名前	意味
\$IP-target_name	インスタンスのクラスの Bayesian 分子 (前のクラスの確率と条件付きのインスタンスの属性値の確率の積)。
\$ILP-target_name	後者の自然対数。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

IBM Data WH Naive Bayes ナゲット - 「設定」 タブ

「設定」タブでは、モデルのスコアリングのオプションを設定することができます。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

レコードのスコアリングのために割り当てられたクラスの確率を計算する. (ディシジョン・ツリーと Naive Bayes のみ) 選択した場合、このオプションは、追加のモデリング・フィールドに確信 (確率) フィールドと予測フィールドが含まれることを意味します。このチェックボックスをオフにすると、予測のみのフィールドが生成されます。

小さなデータ・セットまたはバランスのとれていないデータ・セットの確率の正確度を向上させます。 確率を計算する場合、このオプションは、推定時に 0 の確率を回避するための分推定手法を呼び出します。この種類の確率の推定は、速度が遅くなることがありますが、小さいまたは非常に偏りのあるデータセットによりよい結果を与えることができます。

IBM Data WH KNN モデル・ナゲット

KNN のモデルナゲットは、モデルをスコアリングするためのオプションを設定する手段を提供します。

KNN モデル・ナゲットを含むストリームを実行すると、ノードは1つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 17. KNN のモデル・スコアリング・フィールド	
追加フィールドの名前	意味
\$KNN-target_name	現在のレコードの予測値。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

IBM Data WH KNN ナゲット - 「設定」 タブ

「設定」タブでは、モデルのスコアリングのオプションを設定することができます。

距離測度: 距離が大きいほど、相違が大きくなります。以下のオプションがあります。

- **ユークリッド.** (デフォルト) 2つのポイント間の距離は、直線で結合することによって計算されます。
- **マンハッタン.** 2つのポイント間の距離は、それらの座標間の絶対差の合計として計算されます。
- **キャンベラ.** マンハッタン距離と似ていますが、起点に近いデータ・ポイントにより敏感です。
- **最大.** 2つの点の間の距離は、座標ディメンションに沿った差の最大値として計算されます。

最近傍の数 (k): 特定のケースの最近傍の数。使用する近傍の数が多いほど、より正確なモデルが作成されるとは限らないことに注意してください。

kを選択することにより、オーバーフィット（特に「ノイズの多い」データの場合は重要な場合があります）の防止と解決（同様のインスタンスに対する異なる予測）のバランスを制御します。各データセットごとにkの値（値は1～数十の範囲）を調整する必要があります。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

距離を計算する前に測定を標準化します。 このオプションを選択すると、距離値を計算する前に連続型入力フィールドの測定値が標準化されます。

コア・セットを使用して、大規模なデータ・セットのパフォーマンスを向上させます。 このオプションを選択すると、大規模データ・セットが関係する場合の計算を高速化するために、コア・セット・サンプリングが使用されます。

Netezza 分裂クラスタリングのモデル・ナゲット

分裂クラスタリングは、モデルをスコアリングするためのオプションを設定する手段を提供します。

分裂クラスタリング・モデル・ナゲットを含むストリームを実行すると、ノードは2つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 18. 分裂クラスタリングのモデル・スコアリング・フィールド	
追加フィールドの名前	意味
\$DC-target_name	現在のレコードが割り当てられているサブクラスターの識別子。
\$DCD-target_name	現在のレコードのサブクラスター中心からの距離。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

Netezza 分裂クラスタリング ナゲット - 「設定」 タブ

「設定」タブでは、モデルのスコアリングのオプションを設定することができます。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

距離測度: 距離が大きいほど、相違が大きくなります。以下のオプションがあります。

- **ユークリッド.** (デフォルト) 2つのポイント間の距離は、直線で結合することによって計算されます。
- **マンハッタン.** 2つのポイント間の距離は、それらの座標間の絶対差の合計として計算されます。
- **キャンベラ.** マンハッタン距離と似ていますが、起点に近いデータ・ポイントにより敏感です。
- **最大.** 2つの点の間の距離は、座標ディメンションに沿った差の最大値として計算されます。

適用された階層レベル。データに適用する必要がある階層のレベル。

IBM Data WH PCA モデル・ナゲット

PCA のモデルナゲットは、モデルをスコアリングするためのオプションを設定する手段を提供します。

PCA モデル・ナゲットを含むストリームを実行すると、デフォルトでは、ノードは1つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 19. PCA のモデル・スコアリング・フィールド	
追加フィールドの名前	意味
\$F-target_name	現在のレコードの予測値。

主成分の数 (N) ... に 1 より大きい値を指定した場合 モデル作成ノードまたはモデル・ナゲットのいずれかのフィールドを選択してストリームを実行すると、ノードは各コンポーネントに新しいフィールドを追加します。この場合、フィールド名には -n という接尾辞が付加されます。n は成分の数です。例えば、pca という名前で 3 つの成分を含むモデルの場合、新規フィールド名は \$F-pca-1、\$F-pca-2、および \$F-pca-3 になります。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

注: 小文字のフィールド名が使用されている場合、モデルのスコアリング時にエラーが発生することがあります。これは、Db2 Data Warehouse の既知の不具合です。回避策としては、スコアリング前にすべてのフィールドを大文字に名前変更してください。

IBM Data WH PCA ナゲット - 「設定」 タブ

「設定」タブでは、モデルのスコアリングのオプションを設定することができます。

投影に使用する主成分の数。 データ・セットの削減先となる基本コンポーネントの数。この値は、属性（入力フィールド）の数を超えてはなりません。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

Netezza 回帰ツリー・モデル・ナゲット

回帰ツリーのモデルナゲットは、モデルをスコアリングするためのオプションを設定する手段を提供します。

回帰ツリー・モデル・ナゲットを含むストリームを実行すると、デフォルトでは、ノードは1つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 20. 回帰ツリーのモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$I-target_name	現在のレコードの予測値。

モデル作成ノードまたはモデル・ナゲットで「**推定分散を計算**」を選択してストリームを実行した場合、さらにフィールドを追加します。

表 21. 回帰ツリーのモデルスコアリング・フィールド - 追加	
追加フィールドの名前	意味
\$IV-target_name	予測値の推定分散。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

Netezza 回帰ツリー・ナゲット - 「モデル」 タブ

「モデル」タブには、回帰ツリー・モデルの「予測値の重要度」がグラフィカル形式で表示されます。棒の長さは、予測値の重要度を表しています。

注：IBM Netezza Analytics バージョン 2.x 以前で作業している場合、回帰ツリー・モデルの内容はテキスト形式でのみ表示されます。

これらのバージョンの場合、以下の情報が表示されます。

- テキストの各行は、ノードまたは葉に対応しています。
- インデントはツリー・レベルを反映します。
- ノードの場合、分割の条件が表示されます。
- 葉の場合、割り当てられたクラスのラベルが表示されます。

Netezza 回帰ツリー - 「設定」 タブ

「設定」タブでは、モデルのスコアリングのオプションを設定することができます。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

推定分散を計算します。割り当てられたクラスの差異を出力に含めるかどうかを示します。

Netezza 回帰ツリー ナゲット - 「ビューア」 タブ

「ビューア」タブには、SPSS Modeler で回帰ツリー モデルが表示されるのと同じ方法で、ツリー モデルのツリー プレゼンテーションが表示されます。

注：モデルが IBM Netezza Analytics バージョン 2.x 以前で構築されている場合、「ビューア」タブは空です。

IBM Data WH 線型モデル・ナゲット

線型モデルナゲットは、モデルをスコアリングするためのオプションを設定する手段を提供します。

線型モデル・ナゲットを含むストリームを実行すると、ノードは1つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 22. 線型のモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$LR-target_name	現在のレコードの予測値。

IBM Data WH 線型ナゲット - 「設定」 タブ

「設定」タブでは、モデルのスコアリングのオプションを設定することができます。

入力フィールドを含める: 選択すると、元の入力フィールドをすべて下流に渡し、追加のモデリング・フィールドをデータの各行に添付します。このボックスをオフにすると、レコード ID フィールドと追加のモデリング・フィールドのみが渡されるため、ストリームがより早く実行されます。

Netezza 時系列モデル・ナゲット

モデル・ナゲットは、時系列モデリング操作の出力へのアクセスを提供します。出力は以下のフィールドから構成されています。

表 23. 時系列のモデル出力フィールド	
フィールド	説明
TSID	時系列の識別子。モデル作成ノードの「フィールド」タブの「時系列 ID」で指定したフィールドの内容。詳しくは、トピック 65 ページの『 Netezza 時系列フィールド・オプション 』を参照してください。
時刻	現在の時系列の時間。
HISTORY	過去のデータ値 (予測に使用された値)。モデル・ナゲットの「設定」タブでオプション「 結果の過去の値を含める 」を選択した場合にのみ、このフィールドが含まれます。
\$TS-INTERPOLATED	使用される補間の値。モデル・ナゲットの「設定」タブでオプション「 結果に補間値を含める 」を選択した場合にのみ、このフィールドが含まれます。補間は、モデル作成ノードの「作成オプション」タブのオプションです。
\$TS-FORECAST	時系列の予測値。

モデル出力を表示するには、テーブル・ノード (ノード・パレットの「出力」タブ) をモデル・ナゲットに接続してテーブル・ノードを実行します。

Netezza 時系列ナゲット - 「設定」 タブ

「設定」タブで、モデル出力をカスタマイズするオプションを指定できます。

モデル名。 モデル作成ノードの「モデル・オプション」タブで指定した、モデルの名前。

その他のオプションは、モデル作成ノードの「モデル・オプション」タブと同じです。

IBM Data WH 一般化線型モデル・ナゲット

モデル・ナゲットは、モデリング操作の出力へのアクセスを提供します。

一般化線型モデル・ナゲットを含むストリームを実行すると、ノードは 1 つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 24. 一般化線形のモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$GLM-target_name	現在のレコードの予測値。

「モデル」タブには、モデルに関連するさまざまな統計が表示されます。

出力は以下のフィールドから構成されています。

表 25. 一般化線形の出力フィールド	
出力フィールド	説明
パラメーター	モデルで使用されるパラメーター (予測変数)。切片 (回帰モデルの定数項) のほか、数値型および名義型の列です。
ベータ	相関係数 (モデルの線型要素)。
標準誤差	データの標準誤差。
テスト	パラメーターの妥当性の評価に使用される検定統計。
p-値	パラメーターが有意であると想定する場合のエラーの確率。
残差要約	

表 25. 一般化線形の出力フィールド (続き)	
出力フィールド	説明
残差のタイプ	要約値が表示される予測の残差の種類。
RSS	残差の値。
df	残差の自由度。
p-値	エラーの確率。値が大きいほど適合度の低いモデルを示します。値が小さいほど適合度が高くなります。

IBM Data WH 一般化線形モデル・ナゲット - 「設定」 タブ

「設定」タブで、モデル出力をカスタマイズできます。

オプションは、モデル作成ノードの「スコアリング・オプション」に表示されているものと同じです。詳しくは、トピック 56 ページの『[IBM Data WH 一般化線形モデル・オプション - スコアリング・オプション](#)』を参照してください。

IBM Data WH TwoStep モデル・ナゲット

TwoStep モデル・ナゲットを含むストリームを実行すると、そのレコードの所属クラスターと割り当てられたクラスターの中心からの距離を含む 2 つの新規フィールドが追加されます。\$TS-Twostep という名前の新規フィールドは所属クラスター用で、\$TSP-Twostep という名前の新規フィールドはクラスターの中心からの距離用です。

IBM Data WH TwoStep ナゲット - 「モデル」 タブ

「モデル」タブには、クラスターのフィールドの要約統計量および分布を表示する各種のグラフィック表示があります。モデルからデータをエクスポートしたり、ビューをグラフィックとしてエクスポートしたりすることができます。

第 6 章 IBM Db2 for z/OS によるデータベース モデリング

IBM SPSS Modeler および IBM Db2 for z/OS

SPSS Modeler では、Db2 for z/OS との統合をサポートしており、Db2 for z/OS サーバーで高度な分析を実行する機能を提供します。これらの機能には、SPSS Modeler グラフィカル ユーザー インターフェースおよびワークフロー指向の開発環境でアクセスできます。この方法により、Db2 for z/OS 環境で直接、IBM Db2 Analytics Accelerator を活用してデータ マイニング アルゴリズムを実行できます。

SPSS Modeler は、Db2 for z/OS の次のアルゴリズムの統合をサポートします。

- ディジション・ツリー
- K-Means
- Naive Bayes
- 回帰ツリー
- TwoStep

IBM Db2 for z/OS との統合の要件

Db2® for z/OS® および IBM Db2 Analytics Accelerator for z/OS を使用してデータベース内モデル作成を実行する場合、以下の条件が前提条件となります。場合によっては、これらの条件が満たされているかデータベース管理者に問い合わせて確認してください。サポートされるバージョンを含む、要件の詳細については、[Software Product Compatibility Reports](#) を参照してください。

- ローカル モードで動作している IBM SPSS Modeler、あるいは Windows または UNIX 上の SPSS Modeler Server インストール済み環境に対して動作している SPSS Modeler
- Db2 Analytics Accelerator for z/OS と併用する Db2 for z/OS
- IBM SPSS Data Access Pack
- SPSS Modeler Server を実行するサーバーでは、以下のいずれかのシステム：
 - IBM Db2 Data Server Driver for ODBC and CLI
 - Db2 for z/OS 用に構成された ODBC データ・ソースを使用する任意のバージョンの Db2 for Linux®, UNIX, and Windows
- Db2 Connect for System z® のライセンス
- SPSS Modeler で有効化された SQL の生成および最適化
- Db2 z/OS データベース内マイニングには、アクセラレーター専用テーブル (AOT) またはアクセラレーテッド・テーブル、および INZA のサポートが必要です。IDAA INZA は IDAA 5.1 で導入されました。つまり、Db2 z/OS データベース内マイニング・ノードは、前のバージョンの IDAA では動作しません。

IDAA 対応の DSN を Modeler で使用する場合、その DSN を使用してデータベース入力ノードに返されるテーブルのリストに表示される専用テーブルは、AOT またはアクセラレーテッド・テーブルです。

IBM Db2 Analytics Accelerator for z/OS との統合の有効化

Db2 Analytics Accelerator for z/OS との統合を有効にする手順は、次のとおりです。

- Db2 for z/OS および Db2 Analytics Accelerator for z/OS の構成
- ODBC ソースの作成
- IBM SPSS Modeler での IBM Db2 for z/OS の統合の有効化
- SPSS Modeler での SQL の生成と最適化を有効にする

- IBM SPSS Modeler Server Scoring Adapter for Db2 for z/OS の有効化
- IBM SPSS Modeler での IBM Db2 クライアントを使用した DSN の構成

IBM Db2 for z/OS および IBM Analytics Accelerator for z/OS の構成

Db2 for z/OS および Analytics Accelerator for z/OS の構成方法は、次の Web サイトで説明されています。

[DB2 Analytics Accelerator for z/OS](#)

IBM Db2 for z/OS および IBM Db2 Analytics Accelerator の ODBC ソースの作成

Db2 for z/OS と IBM Db2 Analytics Accelerator の間の接続を有効にする方法については、次の Web サイトを参照してください。

- バージョン 4 の場合: [Db2 Analytics Accelerator for z/OS 4.1.0](#)
- バージョン 3 の場合: [Db2 Analytics Accelerator for z/OS 3.1.0](#)
- [Enabling query acceleration with IBM DB2 Analytics Accelerator for ODBC and JDBC applications without modifying the applications](#)
- [SQL error from ODBC driver when running a query in DB2 Analytics Accelerator for z/OS](#)

IBM SPSS Modeler での IBM Db2 for z/OS の統合の有効化

SPSS Modeler で Db2 for z/OS の統合を有効にするには、以下のステップを実行します。

1. SPSS Modeler config ディレクトリーから `odbc-db2-accelerator-names.cfg` ファイルを開きます。

このファイルが存在しない場合は、作成する必要があります。

2. すべてのデータ ソースの名前とすべてのアクセラレータの名前を追加します。以下に例を示します。

```
dsn1, acceleratorname1
dsn2, acceleratorname2
```

3. アクセラレータ専用テーブル (AOT) のデフォルト CCSID は Unicode です。これを上書きするには、エンコード文字列をアクセラレータ名に追加して、エントリーを変更します。以下に例を示します。

```
dsn1, acceleratorname1, EBCDIC
dsn2, acceleratorname2, UNICODE
```

4. `odbc-db2-accelerator-names.cfg` ファイルを保存して閉じた後、同じディレクトリーから `odbc-db2-custom-properties.cfg` ファイルを開きます。
5. SPSS Modeler は、SQL を使用して IDAA レジスターを設定します。必要な場合は、SQL を目的の値に変更することにより、これらのエントリーを上書きできます。以下に例を示します。

```
current_query_sql_acc, "SET CURRENT QUERY ACCELERATION = ELIGIBLE"
current_get_archive_acc, "SET CURRENT GET_ACCEL_ARCHIVE = NO"
```

6. デフォルトでは、SPSS Modeler は、SQL を使用して、データベース・キャッシュ用の一時テーブルを作成します。必要な場合は、所定のデータベース名を指定して、これを上書きできます。以下に例を示します。

```
[OSZ]
table_create_temp_sql_acc, 'CREATE TABLE <table-name> <(table-columns)> IN DATABASE
NAME_OF_DATABASE_FOR_AOT'
```

7. デフォルトでは、SPSS Modeler は、ODBC 入力ノードで作成された SQL クエリーは再生不可であると判断します。つまり、クエリーを複数回実行すると、異なる結果が返されると判断します。ただし、一部のシナリオでは、これにより Modeler が下流ノードの SQL を生成できなくなる可能性があり、関連する値を Y に変更することでオーバーライドできます。以下に例を示します。

```
assume_custom_sql_replayable, Y
```

8. SPSS Modeler のメインメニューから、「ツール」 > 「オプション」 > 「ヘルパー アプリケーション」をクリックします。
9. 「**IBM DB2 for z/OS**」タブをクリックします。
10. 「**IBM DB2 for z/OS Data Mining 統合を有効化**」を選択して、「**OK**」をクリックします。

注: IDAA テーブルと非 IDAA テーブルを Modeler 内に同時に表示することはできません。

SQL の生成と最適化を有効にする

非常に大きいデータ・セットを扱う確率が高いため、パフォーマンス上の理由により、IBM SPSS Modeler で SQL 生成および最適化を有効にする必要があります。

SPSS Modeler を構成するには、以下のステップを実行します。

1. IBM SPSS Modeler メニューから「ツール」 > 「ストリームのプロパティ」 > 「オプション」を選択します。
2. ナビゲーション ペインの「最適化」オプションをクリックします。
3. 「SQL 生成」オプションが有効になっていること確認します。この設定は、データベースのモデル作成が機能するために必要です。
4. 「SQL 生成の最適化」と「その他の実行を最適化」を選択します (絶対に必要な訳ではありませんが、最適化されたパフォーマンスを得るために、選択することを強くお勧めします)。

IBM SPSS Modeler での IBM Db2 クライアントを使用した DSN の構成

SPSS Modeler で、Db2 用の Db2 クライアントを使用してデータ・ソース名 (DSN) を構成する必要がある場合は、次の手順を実行します。

1. Modeler Server がインストールされているオペレーティング・システムに Db2 クライアントをまだインストールしていない場合は、インストールします。
2. **db2 catalog** コマンドを使用してデータベースをカタログし、Db2 クライアントの db2cli.ini ファイルに新しいデータ・ソースを追加します。定義済みのデータベースのエイリアスを指すようにしてください。
3. データ・アクセスを構成します。詳細な手順については、Modeler の資料を参照してください。
詳しくは、「Modeler Server 管理およびパフォーマンス・ガイド」(ModelerServerAdminPerformance.pdf) のトピック『サーバーアーキテクチャーとハードウェアに関する推奨事項』 > 『データへのアクセス』を参照してください。
4. 手順 2 で定義したデータベースのエイリアスを参照して、odbc.ini に新しい ODBC データ・ソースを作成します。
5. Linux ユーザーまたは UNIX ユーザーの場合:
 - a. ドライバー・ライブラリー libdb2o.so が (libdb2.so の代わりに) 使用されていること、および新しいデータ・ソースに対して 'DriverUnicodeType=1' が定義されていることを確認してください。
 - b. IBM SPSS Data Access Pack のインストール済み環境で、Db2 クライアントのライブラリー・パスが odbc.sh に追加されていることを確認します。
 - c. Modeler Server で、エンコード方式が UTF-16 である ODBC ドライバー・ラッパー・ライブラリー (これは 'libspssodbc_datadirect_utf16.so' と呼ばれます) が使用されていることを確認します。
6. Db2 に接続するユーザーが、次のクエリーの実行に必要な権限を備えていることを確認します。

```
SELECT ACCELERATORNAME FROM SYSACCEL.SYSACCELERATORS
```

IBM Db2 for z/OS でのモデル構築

サポートされているアルゴリズムのそれぞれに、対応するモデル作成ノードがあります。ノード・パレットの「DB モデル作成」タブから、Db2 for z/OS モデル作成ノードにアクセスできます。

データの考慮事項

データ・ソース内のフィールドには、モデル作成ノードに応じてさまざまなデータ型の変数を含めることができます。SPSS Modeler では、データ型は測定の尺度とも呼ばれます。モデル作成ノードの「フィールド」タブでは、入力フィールドと対象フィールドに許可されている測定の尺度の種類を示すアイコンを使用しています。

対象フィールド: 対象フィールドは、値を予測しようとしているフィールドです。対象を指定できる場合、ソース・データ・フィールドのうち 1 つだけを対象フィールドとして選択できます。

「レコード ID」フィールド: 各ケースを一意的に識別するために使用するフィールドを指定します。例えば、これは *CustomerID* などの ID フィールドです。ソースデータに ID フィールドが含まれていない場合は、次の手順に示すように、フィールド作成ノードを使用してこのフィールドを作成することができます。

1. 入力ノードを選択します。
2. ノード・パレットの「フィールド設定」タブから、フィールド作成ノードをダブルクリックします。
3. 領域内のアイコンをダブルクリックしてフィールド作成ノードを開きます。
4. 「派生フィールド」フィールドで、例えば ID を入力します。
5. 「CLEM 式」フィールドで、@INDEX と入力して「OK」をクリックします。
6. フィールド作成ノードをストリームの残りに接続します。

ヌル値の処理

入力データにヌル値が含まれている場合、いくつかの Db2 for z/OS ノードを使用することによって、エラー・メッセージが表示されたり、ストリームの実行時間が長くなる可能性があるため、ヌル値を含むレコードは削除することをお勧めします。以下の方法を使用します。

1. 条件抽出ノードを入力ノードに接続します。
2. 条件抽出ノードの「モード」を「破棄」に設定します。
3. 「条件」フィールドに以下を入力します。

```
@NULL(field1) [or @NULL(field2)[... or @NULL(fieldN)]]
```

すべての入力フィールドを含むようにします。

4. 条件抽出ノードをストリームの残りに接続します。

モデル出力

Db2 for z/OS モデル作成ノードを含むストリームは、実行のたびにわずかに異なる結果を生成する可能性があります。それは、データがモデル作成の前に一時テーブルに読み込まれるため、ノードがソース・データを読み取る順序が必ずしも同じではないからです。ただし、この効果によって生まれる相違点は無視できます。

一般的なコメント

- SPSS Collaboration and Deployment Services では、Db2 for z/OS モデル作成ノードを含むストリームを使用してスコアリング設定を作成することはできません。
- Db2 for z/OS ノードで作成されたモデルの場合、PMML エクスポートまたはインポートを行うことはできません。

IBM Db2 for z/OS モデル - フィールド オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

「すべて」ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定の尺度のボタンをクリックして、その測定の尺度のすべてのフィールドを選択します。

target. 1つのフィールドを予測の対象として選択します。一般化線形モデルの場合、この画面の「試行回数」フィールドも参照してください。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1つ以上のフィールドを予測の入力として選択します。

IBM Db2 for z/OS モデル - サーバー・オプション

「サーバー」タブで、モデルが構築される Db2 for z/OS システムを指定します。

- **上流の接続を使用: (デフォルト)** データベース入力ノードなど、上流のノードで指定した接続の詳細を使用します。注: このオプションは、すべての上流ノードが SQL プッシュバックを使用できる場合にのみ有効です。この場合、SQL が完全にすべての上流ノードを実装するため、データベース外のデータを移動する必要はありません。
- **接続へのデータの移動.** ここで指定したデータベースにデータを移動します。データを移動することにより、データが別の IBM データベース、他のベンダーのデータベースにある場合、またはデータがフラット・ファイルにある場合でもモデリングを実行できるようにします。また、ノードが SQL プッシュバックを実行していないためデータが抽出されていない場合、データはここで指定されたデータベースに戻されます。接続を参照して選択するには、「編集」ボタンをクリックします。

注: ODBC データ・ソース名は、各 SPSS Modeler ストリームに効果的に埋め込まれます。あるホスト上で作成されたストリームが別のホスト上で実行された場合、データ・ソースの名前はそれぞれのホストで同じである必要があります。また、各入力ノードまたはモデル作成ノードで、「サーバー」タブから異なるデータ・ソースを選択できます。

IBM Db2 for z/OS モデル - モデル・オプション

「モデル・オプション」タブで、モデルの名前を指定するか、自動的に名前を生成するかを選択できます。

モデル名 ターゲットまたは ID フィールド (その指定がない場合はモデル タイプ) に基づいてモデル名を生成、またはカスタム名を指定することができます。

名前が使用されている場合は既存モデルを置換。 このチェック ボックスを選択する場合、同じ名前の既存モデルは上書きされます。

IBM Db2 for z/OS モデル - K-Means

K-Means ノードは、クラスター分析の方法を提供する *k-means* アルゴリズムを実行します。このノードを使用して、データ・セットをグループにクラスター化できます。

アルゴリズムは距離に基づくアルゴリズムで、距離メトリック (関数) に依存してデータ・ポイント間の類似性を測定します。データ・ポイントは、使用する距離メトリックに従って最も近いクラスターに割り当てられます。

アルゴリズムは、同じ基本プロセスを何度か反復することによって実行します。学習インスタンスは最も近いクラスターに割り当てられます (指定された距離関数に関しては、インスタンスとクラスター中心に適

用されます)。すべてのクラスター中心は、特定のクラスターに割り当てられたインスタンスの平均属性値のベクトルとして再計算されます。

IBM Db2 for z/OS モデル - K-Means のフィールド オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

「すべて」ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定の尺度のボタンをクリックして、その測定の尺度のすべてのフィールドを選択します。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

IBM Db2 for z/OS モデル - K-Means の作成オプション

作成オプションを設定することにより、用途に合わせてモデルの作成をカスタマイズできます。

デフォルト・オプションを使用してモデルを作成する場合は、「実行」をクリックします。

距離尺度: このパラメーターは、データ・ポイント間の距離の測定方法を定義します。距離が大きくなると、非類似度も大きくなります。次のオプションのいずれかを選択します:

- **ユークリッド。** ユークリッド測定は、2 つのデータ・ポイント間の直線距離です。
- **正規化ユークリッド (Normalized Euclidean)。** 正規化ユークリッド測定はユークリッド測定に類似していますが、平方標準偏差によって正規化されます。ユークリッド測定とは異なり、正規化ユークリッド測定はスケール不変でもあります。

クラスター数: このパラメーターは、作成するクラスターの数进行を定義します。

反復の最大数。 アルゴリズムは、同じプロセスを何度か反復します。このパラメーターは、モデルの学習を停止する前の反復数を定義します。

統計。 このパラメーターは、モデルに含まれる統計の数を定義します。次のオプションのいずれかを選択します:

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注: このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「なし」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし:** モデルのスコアリングに必要な統計のみが含まれます。

結果を再現。 ランダム・シードを設定して分析を複製する場合は、このチェック・ボックスを選択します。整数を指定するか、「生成」をクリックして整数の疑似乱数を作成します。

IBM Db2 for z/OS モデル - Naive Bayes

Naive Bayes は、分類の問題に対応する有名なアルゴリズムです。提示されたすべての予測変数は相互に依存関係がないものとして処理されるので、モデルは naïve と命名されます。Naive Bayes は拡張性のある高速のアルゴリズムであり、複数の属性と対象属性の組み合わせに対して、条件に関連する確率を計算します。学習データから、個別の確率が計算されます。各入力変数の各値カテゴリーを計算単位とすると、この確率は各対象クラスの確率を表します。

IBM Db2 for z/OS モデル - ディシジョン・ツリー

ディシジョン・ツリーは、分類モデルを示す階層構造です。ディシジョン・ツリー・モデルによって、学習データのセットから将来の観測値を予測または分類する分類システムを開発できます。分類は、ブランチが分類の分割点を示すツリー構造の形式です。分割は、停止点に達するまでデータをサブグループに繰り返し分割します。停止点のツリーノードは、葉と呼ばれます。各葉は、クラス・ラベルというラベルを、サブグループのメンバー、またはクラスに割り当てます。

IBM Db2 for z/OS モデル - ディシジョン ツリーのフィールド オプション

「フィールド」タブで、上流のノードですでに定義されているフィールドの役割設定を使用するか、手動でフィールドの割り当てを行うかを選択します。

定義済みの役割を使用します。 上流のデータ型ノード (または上流の入力ノードの「データ型」タブ) の役割設定 (対象、予測など) を使用します。

カスタム・フィールド割り当てを使用します。 この画面で対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印ボタンを使用して、このリストの項目を画面右側のさまざまな役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

「すべて」ボタンをクリックしてリスト内のすべてのフィールドを選択するか、各測定の尺度のボタンをクリックして、その測定の尺度のすべてのフィールドを選択します。

target. 1 つのフィールドを予測の対象として選択します。

レコード ID: 一意のレコード ID として使用されるフィールド。このフィールドの値は各レコードで一意である必要があります (例: カスタマー ID 番号)。

インスタンスの重み。 ここでフィールドを指定すると、デフォルトのクラスの重み (対象フィールドのカテゴリごとの重み) の代わりに、またはデフォルトのクラスの重みに加えて、インスタンスの重み (入力データの行ごとの重み) を使用できます。ここで指定するフィールドは、入力データの各行の数値の重みを含むフィールドでなければなりません。

予測値 (入力): 1 つ以上の入力フィールドを選択します。これは、データ型ノードのフィールドの役割を「入力」に設定する場合と類似しています。

IBM Db2 for z/OS モデル - ディシジョン ツリーの作成オプション

ツリーの成長には、以下の作成オプションを使用できます。

成長の測定。 これらのオプションは、ツリーの成長を測定する方法を制御します。

- **不純度の測定:** この測定では、ツリーを分割するのに最適な場所が評価されます。これは、サブグループまたはデータのセグメントにおける変動性の測定です。不純度の測定値が低い場合は、グループ内のほとんどのメンバーの基準フィールドまたは対象フィールドの値が類似していることを示します。

サポートされる測定は、「エントロピー」と「Gini」です。これらの測定は、ブランチの所属カテゴリーの確率に基づいています。

- **最大ツリー深さ。** ルート・ノード下でツリーが成長可能な最大レベル数 (サンプルが再帰的に分割される回数)。このプロパティのデフォルト値は 10 であり、このプロパティに設定できる最大値は 62 です。

注: モデル・ナゲットのビューアーにモデルがテキスト表示される場合、最大 12 のツリー・レベルが表示されます。

分割の基準。 これらのオプションは、ツリーの分割をいつ停止するかを制御します。

- **分割の改善度の最小変化量。** 新しい分割がツリーで作成される前に不純度を減少する必要のある最小数。ツリー構築の目標は、似かよった出力値を持つサブグループを作成して、それぞれのノード内における不純度を最小にすることです。ブランチが適切に分割されて不純度が分割基準によって指定された値を下回ると、ブランチは分割されません。

- **分割の最小インスタンス数。** 分割可能な最小レコード数。分割されていないレコードがこの数より少ない場合、これ以上分割は行われません。このフィールドを使用すると、ツリー内に小さいサブグループが作成されないようにすることができます。

統計。 このパラメーターは、モデルに含まれる統計の数を定義します。次のオプションのいずれかを選択します:

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注: このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「なし」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし:** モデルのスコアリングに必要な統計のみが含まれます。

IBM Db2 for z/OS モデル - ディシジョン ツリー ノード - クラスの重み

ここでは、重みを各クラスに割り当てることができます。デフォルトでは、1の値をすべてのクラスに割り当て、重みを等しくします。異なるクラス・ラベルに異なる数値の重みを指定することによって、アルゴリズムが特定のクラスの学習セットに重みを付けるよう指定します。

重みを変更するには、「**重み**」列で変更する重みをダブルクリックし、必要に応じて変更します。

値。 対象フィールドの可能な値から派生した、クラス ラベルのセット。

重み。 特定のクラスに割り当てられる重み。大きな重みをクラスに割り当てると、モデルは、他のクラスと比べてそのクラスに対して敏感になります。

インスタンスの重みと組み合わせてクラスも重みを使用できます。

IBM Db2 for z/OS モデル - ディシジョン ツリー ノード - ツリーの剪定

剪定オプションを使用して、ディシジョン・ツリーの剪定基準を指定できます。剪定の目的は、新しいデータに対して予期した精度が改善されない成長しすぎたサブグループを削除することによって、オーバーフィットのリスクを軽減することにあります。

剪定の測定。 デフォルトの剪定測定 **正解率**では、ツリーから葉を削除した後、モデルの推定精度が許容限度内に維持されます。剪定を適用しながらクラスの重みを考慮に入れる場合、「**重みつき精度**」を使用します。

剪定するデータ: 新しいデータに対する想定された精度を推定するには、学習データの一部またはすべてを使用できます。また、指定されたテーブルの別の剪定データセットを使用できます。

- **すべての学習データを使用:** このオプション (デフォルト) は、モデルの精度を推定するためにすべての学習データを使用します。
- **剪定に次のパーセンテージ学習データを使用:** 剪定データに指定した割合で、一方は学習用、一方は剪定用です。
- **ストリームを実行するごとにデータを同じ方法で区分するようランダム シードを指定する場合、「結果を再現」を選択します。** 「**剪定に使用するシード**」フィールドで整数を指定するか、または「**生成**」をクリックすると、擬似無作為の整数を作成します。
- **既存のテーブルのデータを使用:** モデルの精度を推定するために個別の剪定データセットのテーブル名を指定します。学習データを使用するより信頼性が高いと見なされます。

IBM Db2 for z/OS モデル - 回帰ツリー

回帰ツリーは、数値型対象フィールドの値に基づいて同じ種類のサブセットを派生させるために、ケースのサンプルを繰り返し分割するツリーベースのアルゴリズムです。ディシジョン・ツリーと同様に、回帰ツリーはツリーの葉が十分に小さいか、均一なサブセットに対応するサブセットにデータを分解します。分割は、葉の平均値で合理的に予測できるよう、対象の属性の値のばらつきを減少させるために選択されます。

IBM Db2 for z/OS モデル - 回帰ツリーの作成オプション - ツリーの成長

ツリーの成長とツリーの剪定の作成オプションを設定できます。

ツリーの成長には、以下の作成オプションを使用できます。

最大ツリー深さ。 ルート・ノード下でツリーが成長可能な最大レベル数 (サンプルが再帰的に分割される回数)。デフォルトは 62 で、モデリングのための最大ツリー深度です。

注: モデル・ナゲットのビューアーにモデルがテキスト表示される場合、最大 12 のツリー・レベルが表示されます。

分割の基準。 これらのオプションは、ツリーの分割をいつ停止するかを制御します。

- **評価指標の分割。** このクラス評価測定では、ツリーの分割に最適な場所が評価されます。

注: 現在、使用可能なオプションは「分散」のみです。

- **分割の改善度の最小変化量。** 新しい分割がツリーで作成される前に不純度を減少する必要がある最小数。ツリー構築の目標は、似かよった出力値を持つサブグループを作成して、それぞれのノード内における不純度を最小にすることです。ブランチが適切に分割されて不純度が分割基準によって指定された値を下回ると、ブランチは分割されません。
- **分割の最小インスタンス数。** 分割可能な最小レコード数。分割されていないレコードがこの数より少ない場合、これ以上分割は行われません。このフィールドを使用すると、ツリー内に小さいサブグループが作成されないようにすることができます。

統計。 このパラメーターは、モデルに含まれる統計の数を定義します。次のオプションのいずれかを選択します:

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注: このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「なし」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし:** モデルのスコアリングに必要な統計のみが含まれます。

IBM Db2 for z/OS モデル - 回帰ツリーの作成オプション - ツリーの剪定

剪定オプションを使用して、回帰ツリーの剪定基準を指定できます。剪定の目的は、新しいデータに対して予期した精度が改善されない成長しすぎたサブグループを削除することによって、オーバーフィットのリスクを軽減することにあります。

プルーニング測定。 剪定の測定により、ツリーから葉を削除した後、モデルの推定精度が許容限度内に維持されます。次のいずれかを選択できます。

- **MSE** 平均二乗誤差-(デフォルト) 適合線がデータ・ポイントにどの程度近いかを測定します。
- **r²** R-squared - 回帰モデルによって説明される従属変数の変動の比率です。
- **Pearson:** Pearson の相関係数 - 正規分布されている線形従属変数間の関係の強さを測定します。
- **Spearman:** Spearman の相関係数 - ピアソンの相関に従って弱いと思われるが、実際には強いと考えられる非線形な関係を検出します。

剪定するデータ: 新しいデータに対する想定された精度を推定するには、学習データの一部またはすべてを使用できます。また、指定されたテーブルの別の剪定データセットを使用できます。

- **すべての学習データを使用:** このオプション (デフォルト) は、モデルの精度を推定するためにすべての学習データを使用します。
- **剪定に次のパーセンテージ学習データを使用:** 剪定データに指定した割合で、一方は学習用、一方は剪定用です。

ストリームを実行するごとにデータを同じ方法で区分するようランダム シードを指定する場合、「結果を再現」を選択します。「剪定に使用するシード」フィールドで整数を指定するか、または「生成」をクリックすると、擬似無作為の整数を作成します。

- **既存のテーブルのデータを使用:** モデルの精度を推定するために個別の剪定データセットのテーブル名を指定します。学習データを使用するより信頼性が高いと見なされます。

IBM Db2 for z/OS モデル - TwoStep

TwoStep ノードは、大規模データセットにわたってデータをクラスタ化する方法を提供する TwoStep アルゴリズムを実装します。

このノードを使用すると、使用可能なリソース (例えば、メモリーや時間の制約) を考慮しながら、データをクラスタ化できます。

TwoStep アルゴリズムは、次の方法でデータをクラスタ化するデータベース マイニング アルゴリズムです。

1. クラスタ機能 (CF) ツリーが作成されます。このバランスに優れたツリーは、類似の入力レコードが同じツリー ノードの一部となる階層クラスタリングのためのクラスタ機能を格納します。
2. CF ツリーの葉は、最終クラスタリングの結果を生成するためにメモリー内で階層的にクラスタ化されます。クラスタの最適な数は自動的に決定されます。クラスタの最大数を指定する場合、指定された制限内のクラスタの最適な数が決定されます。
3. クラスタリングの結果は、K-Means アルゴリズムに似たアルゴリズムがデータに適用される 2 番目のステップで調整されます。

IBM Db2 for z/OS モデル - TwoStep のフィールド オプション

フィールド オプションを設定することにより、上流ノードで定義されているフィールドの役割の設定を指定できます。フィールド割り当てを手動で行うこともできます。

項目の選択: 上流のデータ型ノードまたは上流の入力ノードの「データ型」タブの役割設定を使用するには、このオプションを選択します。役割設定は、例えば対象や予測です。

カスタム・フィールド割り当てを使用します。 対象、予測、およびその他の役割を手動で割り当てる場合、このオプションを選択します。

フィールド。 矢印を使用して、このリストの項目を右側の役割フィールドに手動で割り当てます。アイコンは、各役割フィールドの有効な測定の尺度を示します。

レコード ID: 一意のレコード ID として使用されるフィールド。

予測値 (入力): 1 つ以上のフィールドを予測の入力として選択します。

IBM Db2 for z/OS モデル - TwoStep 作成オプション

作成オプションを設定することにより、用途に合わせてモデルの作成をカスタマイズできます。

デフォルト・オプションを使用してモデルを作成する場合は、「実行」をクリックします。

距離測定: このパラメーターは、データ・ポイント間の距離の測定方法を定義します。距離が大きくなると、非類似度も大きくなります。オプションは以下のとおりです。

- **対数尤度。** この尤度測定により、変数の確率分布を求めます。連続型変数は正規分布しているものと仮定され、カテゴリー変数は多項分布しているものと仮定されます。すべての変数は独立しているものと仮定されます。

クラスター番号。 このパラメーターは、作成するクラスターの数を定義します。以下のオプションがあります。

- **クラスターの数を自動的に計算します。** クラスタの数が自動的に計算されます。「最大」フィールドで最大クラスター数を指定できます。
- **クラスターの数を指定します。** 作成するクラスタの数を指定します。

統計。このパラメーターは、モデルに含まれる統計の数を定義します。以下のオプションがあります。

- **すべて。** 列に関連したすべての統計と値に関連したすべての統計が含まれます。

注：このパラメーターを使用すると統計の最大数が含まれるため、システムのパフォーマンスに影響する場合があります。モデルをグラフィカル形式で表示しない場合は、「なし」を指定します。

- **列。** 列に関連した統計が含まれます。
- **なし。** モデルのスコアリングに必要な統計のみが含まれます。

結果を再現。ランダム・シードを設定して分析を複製する場合は、このチェック・ボックスを選択します。整数を指定するか、「生成」をクリックして整数の疑似乱数を作成します。

IBM Db2 for z/OS モデル - TwoStep ナゲット - 「モデル」タブ

「モデル」タブには、クラスターのフィールドの要約統計量および分布を表示する各種のグラフィック表示があります。モデルからデータをエクスポートしたり、ビューをグラフィックとしてエクスポートしたりすることができます。

IBM Db2 for z/OS モデルの管理

Db2 for z/OS モデルは、他の IBM SPSS Modeler のモデルのようにキャンバスおよび「モデル」パレットに追加され、ほぼ同じように使用できます。

Db2 for z/OS で直接的にデータをスコアリングするには、以下の手順を実行します。

1. データが配置されている Db2 for z/OS データベースに SPSS Scoring Adapter をインストールします。
2. ストリームが、データが配置されている Db2 for z/OS データベースに接続することを確認します。

IBM Db2 for z/OS モデルのスコアリング

モデルは、キャンバス上で金色のモデル・ナゲット・アイコンで示されます。ナゲットの主な目的は、データをスコアリングし、予測を生成、またはモデルのプロパティの詳細な分析を可能にすることです。スコアは、このセクションで後述するように、ナゲットにテーブル・ノードを接続し、ストリームのそのブランチを実行することによって見えるようにすることができます、1つ以上の追加データ・フィールドの形式で追加されます。ディシジョンツリーや回帰ツリーのダイアログ・ボックスなど、いくつかのナゲットのダイアログボックスには、さらにモデルの視覚的表現を提供する「モデル」タブがあります。

追加フィールドは、対象フィールドの名前に追加された接頭辞 \$<id>- によって区別されます。<id> はモデルによって異なり、追加される情報の種類を識別します。さまざまな識別子について、それぞれのモデル・ナゲットのトピックで説明されています。

スコアを表示するには、次の手順を実行します。

1. テーブル・ノードをモデル・ナゲットに接続します。
2. テーブル・ノードを開きます。
3. 「実行 (Run)」をクリックします。
4. テーブル出力ウィンドウの右側にスクロールし、追加フィールドとそれらのスコアを表示します。

注：スコアリング・プロセスは、アクセラレーターでは実行されず、Db2 で実行されます。そのため、スコアリングの入力テーブルは物理的に Db2 に配置する必要があります。したがって、スコアリング入力としては、Db2 ベースのテーブルまたはアクセラレーテッド・テーブルのみを使用できます。ストリームがアクセラレーター専用テーブルを使用する場合、「THE STATEMENT CANNOT BE EXECUTED BY DB2 OR IN THE ACCELERATOR.」というエラーが発生します。

IBM Db2 for z/OS ディシジョン・ツリー・モデル・ナゲット

ディシジョン・ツリー・モデル・ナゲットでは、モデリング操作の出力を表示し、モデルをスコアリングするためのオプションを設定することもできます。

ディシジョン ツリー モデル ナゲットを含むストリームを実行すると、ノードは2つの新しいフィールドを追加します。その名前は、対象に由来します。

表 26. ディシジョン・ツリーのモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$I-target_name	現在のレコードの予測値。
\$IP-target_name	予測の確信度 (0.0 ～ 1.0)。

注: Db2 for z/OS の制約により、列名は切り捨てられる場合があります。

IBM Db2 for z/OS ディシジョン・ツリー・ナゲット - 「モデル」 タブ

「モデル」タブには、ディシジョン・ツリー・モデルの「予測値の重要度」がグラフィカル形式で表示されます。棒の長さは、予測値の重要度を表しています。

IBM Db2 for z/OS ディシジョン ツリー ナゲット - 「ビューア」 タブ

「ビューア」タブには、SPSS Modeler でディシジョン ツリー モデルが表示されるのと同じ方法で、ツリーモデルのツリー プレゼンテーションが表示されます。

IBM Db2 for z/OS K-Means モデル・ナゲット

K-Means モデル・ナゲットには、クラスター化モデルが取得したすべての情報と、学習データと推定プロセスに関する情報が含まれます。

K-Means モデル・ナゲットを含むストリームを実行すると、そのレコードの所属クラスターと割り当てられたクラスターの中心からの距離を含む2つの新規フィールドが追加されます。新しいフィールド名はモデル名から派生し、クラスター・メンバーシップの場合は \$KM-、クラスター中心からの距離の場合は \$KMD- という接頭部が付きます。例えば、モデルの名前が Kmeans の場合、新規フィールド名は \$KM-Kmeans と \$KMD-Kmeans になります。

注: Db2 for z/OS の制約により、列名は切り捨てられる場合があります。

IBM Db2 for z/OS K-Means ナゲット - 「モデル」 タブ

「モデル」タブには、クラスターのフィールドの要約統計量および分布を表示する各種のグラフィック表示があります。モデルからデータをエクスポートしたり、ビューをグラフィックとしてエクスポートしたりすることができます。

IBM Db2 for z/OS Naive Bayes モデル ナゲット

Naive Bayes モデル ナゲットを含むストリームを実行すると、ノードは2つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 27. Naive Bayes のモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$I-target_name	現在のレコードの予測値。
\$IP-target_name	予測の確信度 (0.0 ～ 1.0)。

注: Db2 for z/OS の制約により、列名は切り捨てられる場合があります。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

IBM Db2 for z/OS 回帰ツリー モデル ナゲット

回帰ツリー モデル ナゲットを含むストリームを実行すると、ノードは2つの新しいフィールドを追加します。その名前は、対象名に由来します。

表 28. 回帰ツリーのモデルスコアリング・フィールド	
追加フィールドの名前	意味
\$I-target_name	現在のレコードの予測値。
\$IS-target_name	予測値の推定標準偏差。

注：Db2 for z/OS の制約により、列名は切り捨てられる場合があります。

モデルナゲットにテーブル・ノードを接続して表のノードを実行することにより、追加フィールドを表示することができます。

IBM Db2 for z/OS 回帰ツリー・ナゲット - 「モデル」 タブ

「モデル」タブには、回帰ツリー・モデルの「予測値の重要度」がグラフィカル形式で表示されます。棒の長さは、予測値の重要度を表しています。

IBM Db2 for z/OS 回帰ツリー ナゲット - 「ビューア」 タブ

「ビューア」タブには、SPSS Modeler で回帰ツリー モデルが表示されるのと同じ方法で、ツリー モデルのツリー プレゼンテーションが表示されます。

IBM Db2 for z/OS TwoStep モデル・ナゲット

TwoStep モデル・ナゲットを含むストリームを実行すると、そのレコードの所属クラスターと割り当てられたクラスターの中心からの距離を含む2つの新規フィールドが追加されます。新規フィールド名はモデル名から派生し、所属クラスターのフィールドには接頭辞の \$TS-、クラスターの中心からの距離のフィールドには接頭辞の \$TSD- が付けられます。例えば、モデルの名前が MDL の場合、新規フィールド名は \$TS-MDL と \$TSD-MDL になります。

特記事項

本書は米国 IBM が提供する製品およびサービスについて作成したものです。この資料の他の言語版を IBM から入手できる場合があります。ただし、これを入手するには、本製品または当該言語版製品を所有している必要がある場合があります。

IBM は、このドキュメントで説明されている製品、サービス、または機能を、他の国では提供していない場合があります。日本で利用可能な製品、サービス、および機能については、日本 IBM の営業担当員にお尋ねください。本書で IBM 製品、プログラム、またはサービスに言及していても、その IBM 製品、プログラム、またはサービスのみが使用可能であることを意味するものではありません。これらに代えて、IBM の知的所有権を侵害することのない、機能的に同等の製品、プログラム、またはサービスを使用することができます。ただし、IBM 以外の製品とプログラムの操作またはサービスの評価および検証は、お客様の責任で行っていただきます。

IBM は、本書に記載されている内容に関して特許権 (特許出願中のものを含む) を保有している場合があります。本書の提供は、お客様にこれらの特許権について実施権を許諾することを意味するものではありません。実施権についてのお問い合わせは、書面にて下記宛先にお送りください。

〒 103-8510

東京都中央区日本橋箱崎町 19 番 21 号

日本アイ・ビー・エム株式会社

法務・知的財産

US-->

For license inquiries regarding double-byte (DBCS) information, contact the IBM Intellectual Property Department in your country or send inquiries, in writing, to:

<!--Intellectual Property Licensing

<!-- Legal and Intellectual Property Law -->

日本 アイ・ビー・エム株式会社

103-8510

Tokyo 103-8510, Japan-->

IBM およびその直接または間接の子会社は、本書を特定物として現存するままの状態を提供し、商品性の保証、特定目的適合性の保証および法律上の瑕疵担保責任を含むすべての明示もしくは黙示の保証責任を負わないものとします。国または地域によっては、法律の強行規定により、保証責任の制限が禁じられる場合、強行規定の制限を受けるものとします。

この情報には、技術的に不適切な記述や誤植を含む場合があります。本書は定期的に見直され、必要な変更は本書の次版に組み込まれます。IBM は予告なしに、随時、この文書に記載されている製品またはプログラムに対して、改良または変更を行うことがあります。

本書において IBM 以外の Web サイトに言及している場合がありますが、便宜のため記載しただけであり、決してそれらの Web サイトを推奨するものではありません。それらの Web サイトにある資料は、この IBM 製品の資料の一部ではありません。それらの Web サイトは、お客様の責任でご使用ください。

IBM は、お客様が提供するいかなる情報も、お客様に対してなんら義務も負うことのない、自ら適切と信ずる方法で、使用もしくは配布することができるものとします。

本プログラムのライセンス保持者で、(i) 独自に作成したプログラムとその他のプログラム (本プログラムを含む) との間での情報交換、および (ii) 交換された情報の相互利用を可能にすることを目的として、本プログラムに関する情報を必要とする方は、下記に連絡してください。

〒 103-8510

東京都中央区日本橋箱崎町 19 番 21 号

日本アイ・ビー・エム株式会社

法務・知的財産

US-->

本プログラムに関する上記の情報は、適切な使用条件の下で使用することができますが、有償の場合もあります。

本書で説明されているライセンス・プログラムまたはその他のライセンス資料は、IBM 所定のプログラム契約の契約条項、IBM プログラムのご使用条件、またはそれと同等の条項に基づいて、IBM より提供されます。

記載されている性能データとお客様事例は、例として示す目的でのみ提供されています。実際の結果は特定の構成や稼働条件によって異なります。

IBM 以外の製品に関する情報は、その製品の供給者、出版物、もしくはその他の公に利用可能なソースから入手したものです。IBM はこれらの製品をテストしていないため、IBM 以外の製品に関連するパフォーマンス、互換性、またはその他のクレームの正確性を確認できません。IBM 以外の製品の性能に関する質問は、それらの製品の供給者をお願いします。

IBM の将来の方向性および指針に関する記述は、予告なく変更または撤回される場合があります。これらは目標および目的を提示するものにすぎません。

本書には、日常の業務処理で用いられるデータや報告書の例が含まれています。より具体性を与えるために、それらの例には、個人、企業、ブランド、あるいは製品などの名前が含まれている場合があります。これらの名前はすべて架空のものであり、類似する個人や企業が実在しているとしても、それは偶然にすぎません。

商標

IBM、IBM ロゴ、および ibm.com は、世界の多くの国で登録された International Business Machines Corporation の商標です。世界中の多くの国で登録されています。他の製品名およびサービス名等は、それぞれ IBM または各社の商標である場合があります。現時点での IBM の商標リストについては、<http://www.ibm.com/legal/copytrade.shtml> をご覧ください。

Adobe、Adobe ロゴ、PostScript、PostScript ロゴは、Adobe Systems Incorporated の米国およびその他の国における登録商標または商標です。

インテル、Intel、Intel ロゴ、Intel Inside、Intel Inside ロゴ、Centrino、Intel Centrino ロゴ、Celeron、Xeon、Intel SpeedStep、Itanium、および Pentium は、Intel Corporation または子会社の米国およびその他の国における商標または登録商標です。

Linux は、Linus Torvalds の米国およびその他の国における登録商標です。

Microsoft、Windows、Windows NT および Windows ロゴは、Microsoft Corporation の米国およびその他の国における商標です。

UNIX は The Open Group の米国およびその他の国における登録商標です。

Java およびすべての Java 関連の商標およびロゴは Oracle やその関連会社の米国およびその他の国における商標または登録商標です。

製品資料に関するご使用条件

これらの資料は、以下のご使用条件に同意していただける場合に限りご使用いただけます。

適用範囲

IBM Web サイトの「ご利用条件」に加えて、以下のご使用条件が適用されます。

個人使用

これらの資料は、すべての著作権表示その他の所有権表示をしていただくことを条件に、非商業的な個人による使用目的に限り複製することができます。ただし、IBM の明示的な承諾をえずに、これらの資料またはその一部について、二次的著作物を作成したり、配布（頒布、送信を含む）または表示（上映を含む）することはできません。

商業利用

これらの資料は、すべての著作権表示その他の所有権表示をしていただくことを条件に、お客様の企業内に限り、複製、配布、および表示することができます。ただし、IBM の明示的な承諾をえずにこれらの資料の二次的著作物を作成したり、お客様の企業外で資料またはその一部を複製、配布、または表示することはできません。

権利

ここで明示的に許可されているもの以外に、資料や資料内に含まれる情報、データ、ソフトウェア、またはその他の知的所有権に対するいかなる許可、ライセンス、または権利を明示的にも黙示的にも付与するものではありません。

資料の使用が IBM の利益を損なうと判断された場合や、上記の条件が適切に守られていないと判断された場合、IBM はいつでも自らの判断により、ここで与えた許可を撤回できるものとさせていただきます。

お客様がこの情報をダウンロード、輸出、または再輸出する際には、米国のすべての輸出入 関連法規を含む、すべての関連法規を遵守するものとします。

IBM は、これらの資料の内容についていかなる保証もしません。これらの資料は、特定物として現存するままの状態を提供され、商品性の保証、特定目的適合性の保証および法律上の瑕疵担保責任を含むすべての明示もしくは黙示の保証責任なしで提供されます。

索引

日本語, 数字, 英字, 特殊文字の順に配列されています。
なお, 濁音と半濁音は清音と同等に扱われています。

[ア行]

- アソシエーション・ルール
 - エキスパート・オプション [16](#)
 - サーバー・オプション [15](#)
 - スコアリング - サーバー・オプション [19](#)
 - スコアリング - 要約オプション [19](#)
 - モデル・オプション [15](#)
- アソシエーション・ルール・モデル
 - Microsoft [16](#)
- 値の補間、IBM Netezza Analytics 時系列 [63](#)
- アプリケーションの例 [3](#)
- 一意のフィールド
 - Oracle Adaptive Bayes Network [30](#)
 - Oracle Apriori [34, 38](#)
 - Oracle Data Mining [27](#)
 - Oracle K-Means [36](#)
 - Oracle MDL [39](#)
 - Oracle Naive Bayes [29](#)
 - Oracle NMF [37](#)
 - Oracle O-Cluster [35](#)
 - Oracle サポート・ベクター・マシン [31](#)
- 一般化線型モデル
 - IBM Netezza Analytics [54-56, 77, 78](#)
- 一般化線型モデル (GLM)
 - Oracle Data Mining [32, 33](#)
- インスタンスの重み、Netezza ツリー・モデル [56](#)
- エクスポート
 - Analysis Services モデル [21](#)
- エントロピー不純度測定法 [57](#)

[カ行]

- 回帰ツリー
 - IBM Db2 for z/OS [87, 91](#)
 - IBM Netezza Analytics [51, 52, 75, 76](#)
- ガウス・カーネル
 - Oracle サポート・ベクター・マシン [31](#)
- 鍵
 - モデル・キー [7](#)
- 季節的傾向分解、IBM Netezza Analytics [62](#)
- 距離関数
 - Oracle K-Means [36](#)
- クラスター数
 - Oracle K-Means [36](#)
 - Oracle O-Cluster [35](#)
- クラスタリング
 - エキスパート・オプション [15](#)
 - サーバー・オプション [15](#)
 - スコアリング - サーバー・オプション [19](#)
 - スコアリング - 要約オプション [19](#)
 - モデル・オプション [15](#)
 - IBM Netezza Analytics [74](#)
- クラスの重み、Netezza ツリー・モデル [56](#)

- クラス・ラベル、Netezza ツリー・モデル [56, 85](#)
- 交差検証
 - Oracle Naive Bayes [29](#)
- コスト
 - Oracle [28](#)
- 誤分類コスト
 - Oracle [28](#)

[サ行]

- 最近傍モデル
 - IBM Netezza Analytics [59, 60, 73, 74](#)
- 最小記述長 (MDL)
 - Oracle Data Mining [39](#)
- 作成オプション
 - IBM Db2 for z/OS [84-88](#)
 - IBM Netezza Analytics [51-53, 57-59, 61, 62, 65, 67, 68](#)
- サポート・ベクター・マシン
 - Oracle Data Mining [31](#)
- シーケンス・クラスタリング
 - モデル・オプション [15](#)
- 時系列
 - IBM Netezza Analytics [65, 67](#)
- 時系列 (IBM Netezza Analytics) [62, 76, 77](#)
- 指数平滑法
 - IBM Netezza Analytics [62](#)
- 事前確率
 - Oracle Data Mining [32](#)
- 収束の許容範囲
 - Oracle サポート・ベクター・マシン [31](#)
- 主成分分析モデル
 - IBM Netezza Analytics [68, 69, 75](#)
- スコアリング [6, 69, 89](#)
- スペクトル分析、IBM Netezza Analytics [62](#)
- 正規化方法
 - Oracle K-Means [36](#)
 - Oracle NMF [37](#)
 - Oracle サポート・ベクター・マシン [31](#)
- 設定
 - IBM Db2 for z/OS および IBM Analytics Accelerator for z/OS [80](#)
- 線型
 - エキスパート・オプション [16](#)
 - サーバー・オプション [15](#)
 - スコアリング - サーバー・オプション [19](#)
 - スコアリング - 要約オプション [19](#)
 - モデル・オプション [15](#)
 - IBM Db2 for z/OS [86](#)
 - IBM Netezza Analytics [51, 59, 76](#)
- 線型カーネル
 - Oracle サポート・ベクター・マシン [31](#)
- 剪定された Naive Bayes モデル
 - Oracle Adaptive Bayes Network [30](#)

[タ行]

- 単一型閾値

単一型閾値 (続き)

Oracle Naive Bayes [29](#)

単一機能モデル

Oracle Adaptive Bayes Network [30](#)

探索 [22, 43](#)

ディシジョン・ツリー

エキスパート・オプション [15](#)

サーバー・オプション [15](#)

スコアリング - サーバー・オプション [19](#)

スコアリング - 要約オプション [19](#)

モデル・オプション [15](#)

IBM Db2 for z/OS [85, 86, 89-91](#)

IBM Netezza Analytics [56-58, 70, 71, 76](#)

Microsoft Analysis Services [9, 11, 19](#)

Oracle Data Mining [34, 35](#)

データ区分 [37](#)

データ区分フィールド

選択 [37](#)

データ検査ノード [22, 43](#)

データの正規化

Oracle モデル [42](#)

データ分割

Oracle モデル [42](#)

データベース内モデル作成 [19](#)

データベース・マイニング

構成 [11](#)

最適化オプション [6](#)

データの準備 [6](#)

モデルの構築 [6](#)

例 [21](#)

IBM SPSS Modeler を使用 [5](#)

データベース・モデル作成

IBM Netezza Analytics [46, 48, 50](#)

Oracle [25-28](#)

展開 [23, 44](#)

[ナ行]

ニューラル・ネットワーク

エキスパート・オプション [16](#)

サーバー・オプション [15](#)

スコアリング - サーバー・オプション [19](#)

スコアリング - 要約オプション [19](#)

モデル・オプション [15](#)

ノード

生成 [21](#)

ノードの生成 [21](#)

[ハ行]

評価 [23, 44](#)

標準偏差

Oracle サポート・ベクター・マシン [31](#)

フィールド・オプション

IBM Db2 for z/OS [83-85, 88](#)

IBM Netezza Analytics [50, 53, 57, 61, 62, 65, 68, 69](#)

複雑性ファクタ

Oracle サポート・ベクター・マシン [31](#)

複雑性ペナルティー [15-17](#)

複数機能モデル

Oracle Adaptive Bayes Network [30](#)

不純度の測定

ディシジョン・ツリー [85](#)

不純度の測定 (続き)

Netezza ディシジョン・ツリー [57](#)

不純度メトリック

Oracle Apriori [34](#)

分割基準

Oracle K-Means [36](#)

文書 [3](#)

分裂クラスタリング

IBM Netezza Analytics [52, 53, 74](#)

ペア単位閾値

Oracle Naive Bayes [29](#)

ポート

Oracle 接続 [26](#)

[マ行]

モデル

エクスポート [7](#)

整合性問題 [7](#)

データベース内モデルの構築 [6](#)

データベース内モデルのスコアリング [6](#)

評価 [23, 44](#)

保存 [7](#)

Analysis Services の管理 [13](#)

Netezza の一覧表示 [51](#)

Netezza の管理 [51](#)

Oracle の参照 [29](#)

モデル・オプション

IBM Db2 for z/OS [83](#)

IBM Netezza Analytics [50, 54, 55, 59, 60, 67](#)

モデル作成ノード

データベース内モデル作成 [6, 9, 11, 13, 19](#)

Microsoft Association Rules [13](#)

Microsoft Decision Trees [13](#)

Microsoft Clustering [13](#)

Microsoft Linear Regression [13](#)

Microsoft Logistic Regression [13](#)

Microsoft Naive Bayes [13](#)

Microsoft Neural Network [13](#)

Microsoft Sequence Clustering [13](#)

Microsoft Time Series [13](#)

モデル・ナゲット

IBM Db2 for z/OS [89-91](#)

IBM Netezza Analytics [54, 70-78](#)

[ヤ行]

葉、Netezza ツリー・モデル [56, 85](#)

要件

IBM Db2 for z/OS [79](#)

[ラ行]

例

アプリケーション ガイド [3](#)

概要 [4](#)

データベース・マイニング [21-23, 43](#)

ロジスティック回帰

エキスパート・オプション [16](#)

サーバー・オプション [15](#)

スコアリング - サーバー・オプション [19](#)

スコアリング - 要約オプション [19](#)

モデル・オプション [15](#)

A

- Adaptive Bayes Network
 - Oracle Data Mining [29](#), [30](#)
- Analysis Services
 - ディシジョン・ツリー [21](#)
 - モデルの管理 [13](#)
 - 例 [21](#)
- Apriori
 - Microsoft [16](#)
 - Oracle Data Mining [37](#), [38](#)
- ARIMA モデル
 - IBM Netezza Analytics [62](#), [66](#)
- Attribute Importance (AI)
 - Oracle Data Mining [39](#), [40](#)

B

- Bayesian Network (ベイズ) モデル
 - IBM Netezza Analytics [62](#), [72](#)

D

- database
 - データベース内モデル作成 [6](#), [9](#), [11](#), [13](#), [19](#)
- Db2 for z/OS モデリング
 - IBM Db2 for z/OS [79](#), [82](#), [83](#)
- DSN
 - 設定 [11](#)

E

- epsilon
 - Oracle サポート・ベクター・マシン [31](#)

G

- Gini 不純度測定法 [57](#)

H

- hostname
 - Oracle 接続 [26](#)

I

- IBM
 - モデルの管理 [51](#)
- IBM Db2 for z/OS
 - 回帰ツリー [86](#)
 - 回帰ツリーの作成オプション [87](#)
 - 回帰ツリー・モデル・ナゲット [91](#)
 - ディシジョン・ツリー [85](#)
 - ディシジョン・ツリーの作成オプション [85](#), [86](#)
 - ディシジョン・ツリーのフィールド・オプション [85](#)
 - ディシジョン・ツリー・モデル・ナゲット [89-91](#)
 - フィールド・オプション [83](#)
 - モデル・オプション [83](#)
 - Db2 for z/OS モデルの管理 [89](#)
 - IBM Db2 Analytics Accelerator for z/OS との統合 [79](#)
 - IBM Db2 for z/OS および IBM Analytics Accelerator for z/OS の構成 [80](#)

IBM Db2 for z/OS (続き)

- IBM Db2 for z/OS との統合の要件 [79](#)
- IBM SPSS Modeler との構成 [82](#), [83](#)
- K-Means [83](#)
- K-means の作成オプション [84](#)
- K-means のフィールド・オプション [84](#)
- K-Means モデル・ナゲット [90](#)
- Naive Bayes [84](#)
- Naive Bayes のモデル・ナゲット [90](#)
- TwoStep [88](#)
- TwoStep 作成オプション [88](#)
- TwoStep フィールド オプション [88](#)
- TwoStep モデル・ナゲット [89](#), [91](#)

IBM Netezza Analytics

- 一般化線型 [54](#)
- 一般化線型モデル・ナゲット [54](#), [77](#), [78](#)
- 一般化線型モデルのオプション [54](#), [55](#)
- 回帰ツリー [51](#)
- 回帰ツリーの作成オプション [51](#), [52](#)
- 回帰ツリー・モデル・ナゲット [75](#), [76](#)
- 最近傍 (KNN) [59](#)
- 時系列 [62](#)
- 時系列の構築オプション [65](#), [67](#)
- 時系列のフィールド・オプション [65](#)
- 時系列のモデル・オプション [67](#)
- 時系列モデル・ナゲット [76](#), [77](#)
- 線型 [59](#)
- 線型ノードの作成オプション [59](#)
- 線型モデル・ナゲット [76](#)
- ディシジョン・ツリー [56](#)
- ディシジョン・ツリーの作成オプション [57](#), [58](#)
- ディシジョン・ツリーのフィールド・オプション [57](#)
- ディシジョン・ツリー・モデル・ナゲット [70](#), [71](#), [76](#)
- フィールド・オプション [50](#)
- 分裂クラスタリング [52](#)
- 分裂クラスタリングの作成オプション [53](#)
- 分裂クラスタリングのフィールド・オプション [53](#)
- 分裂クラスタリングのモデル・ナゲット [74](#)
- ベイズ・ネット [62](#)
- モデル・オプション [50](#)
- モデルの管理 [69](#), [70](#)
- Bayes Net のモデル・ナゲット [72](#)
- Bayes ネットワークの作成オプション [62](#)
- Bayes ネットワークのフィールド・オプション [62](#)
- IBM SPSS Modeler との構成 [46](#), [48](#), [50](#)
- K-Means [60](#)
- K-means の作成オプション [61](#)
- K-means のフィールド・オプション [61](#)
- K-Means モデル・ナゲット [71](#), [72](#)
- KNN モデル・オプション [59](#), [60](#)
- KNN モデル・ナゲット [73](#), [74](#)
- Naive Bayes [62](#)
- Naive Bayes のモデル・ナゲット [73](#)
- PCA [68](#)
- PCA 作成オプション [69](#)
- PCA フィールド・オプション [69](#)
- PCA モデル・ナゲット [75](#)
- TwoStep [67](#)
- TwoStep 作成オプション [68](#)
- TwoStep フィールド オプション [68](#)
- TwoStep モデル・ナゲット [78](#)

IBM SPSS Modeler

- データベース・マイニング [5](#)
- 文書 [3](#)

IBM SPSS Modeler Server [1](#)
IBM SPSS Modeler Solution Publisher
Oracle Data Mining モデル [27](#)

K

K-means
IBM Db2 for z/OS [83](#), [84](#)
IBM Netezza Analytics [60](#), [61](#)
Oracle Data Mining [36](#)
K-Means
IBM Db2 for z/OS [90](#)
IBM Netezza Analytics [71](#), [72](#)
KNN モデル
IBM Netezza Analytics [73](#), [74](#)

M

MDL [29](#)
Microsoft
アソシエーション・ルール・モデル作成 [9](#), [11](#), [19](#)
クラスタリング・モデル作成 [9](#), [11](#), [19](#)
シーケンス・クラスタリング [9](#)
線型 [9](#)
線型モデル作成 [11](#), [19](#)
ディシジョン・ツリー・モデル作成 [9](#), [11](#), [19](#)
ニューラル・ネットワーク [9](#)
ニューラル・ネットワーク・モデル作成 [11](#), [19](#)
モデルの管理 [13](#)
ロジスティック回帰 [9](#)
ロジスティック回帰モデル作成 [11](#), [19](#)
Analysis Services [9](#), [11](#), [19](#)
Naive Bayes モデル作成 [9](#), [11](#), [19](#)
Microsoft Analysis Services [20](#), [21](#)
min-max
データの正規化 [31](#), [42](#)
Minimum Description Length [29](#)

N

Naive Bayes
エキスパート・オプション [15](#)
サーバー・オプション [15](#)
スコアリング - サーバー・オプション [19](#)
スコアリング - 要約オプション [19](#)
モデル・オプション [15](#)
IBM Db2 for z/OS [84](#), [90](#)
IBM Netezza Analytics [62](#), [73](#)
Oracle Data Mining [29](#)
Naive Bayes モデル
IBM Netezza Analytics [73](#)
Oracle Adaptive Bayes Network [30](#)
Netezza
モデルの管理 [51](#)
NMF
Oracle Data Mining [36](#), [37](#)

O

O-Cluster
Oracle Data Mining [35](#)
ODBC
設定 [11](#)

ODBC (続き)
IBM Db2 for z/OS 用の構成 [83](#)
IBM Netezza Analytics 用の構成 [46](#), [48](#), [50](#)
Oracle 用の設定 [25-28](#)
SQL Server を設定する [11](#)
ODM。Oracle Data Mining を参照 [25](#)
Oracle Data Miner [42](#)
Oracle Data Mining
一般化線型モデル (GLM) [32](#), [33](#)
誤分類コスト [41](#)
最小記述長 (MDL) [39](#)
サポート・ベクター・マシン [31](#)
整合性検査 [40](#)
ディシジョン・ツリー [34](#), [35](#)
データの準備 [42](#)
モデルの管理 [40](#), [41](#)
例 [43](#), [44](#)
Adaptive Bayes Network [29](#), [30](#)
Apriori [37](#), [38](#)
Attribute Importance (AI) [39](#), [40](#)
IBM SPSS Modeler との構成 [25-28](#)
K-means [36](#)
Naive Bayes [29](#)
NMF [36](#), [37](#)
O-Cluster [35](#)

P

Publisher ノード
Oracle Data Mining モデル [27](#)

S

sequence clustering (Microsoft)
エキスパート・オプション [19](#)
フィールド・オプション [18](#)
SID
Oracle 接続 [26](#)
Solution Publisher
Oracle Data Mining モデル [27](#)
SQL Server
設定 [11](#)
ODBC 接続 [11](#)
SQL 生成 [6](#)
SVM. サポート・ベクター・マシン を参照 [31](#)

T

time series (Microsoft)
エキスパート・オプション [17](#)
設定オプション [17](#)
モデル・オプション [17](#)
tnsnames.ora ファイル [26](#)
TwoStep
IBM Db2 for z/OS [88](#), [89](#), [91](#)
IBM Netezza Analytics [67](#), [68](#), [78](#)

Z

z 得点
データの正規化 [31](#), [42](#)

[特殊文字]

給仕人

Analysis Services を実行する [15](#), [19](#)

