

IBM SPSS Modeler 18.4 Python スクリプト
とオートメーション・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、[469 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

第 1 章 スクリプトとスクリプト言語.....	1
スクリプトの概要.....	1
スクリプトの種類.....	1
ストリーム・スクリプト.....	1
ストリーム・スクリプトの例 : ニューラル・ネットワークの学習.....	3
Jython コードのサイズ制限.....	3
スタンドアロン スクリプト.....	4
スタンドアロン スクリプトの例 : モデルの保存とロード.....	4
スタンドアロン スクリプトの例 : 特徴量選択モデルの生成.....	4
スーパーノード・スクリプト.....	5
スーパーノード・スクリプトの例.....	6
ストリームでのループと条件付き実行.....	6
ストリームでのループ.....	7
ストリームでの条件付き実行.....	10
スクリプトの実行と中断.....	11
検索と置換.....	11
第 2 章 スクリプト言語.....	15
スクリプト言語の概要.....	15
Python と Jython.....	15
Python スクリプト.....	15
操作.....	16
リスト.....	16
ストリング.....	17
注釈.....	19
ステートメントの構文.....	19
識別子.....	19
コードのブロック.....	20
スクリプトへの引数の引き渡し.....	20
例.....	21
数学メソッド.....	21
非 ASCII 文字の使用.....	23
オブジェクト指向プログラミング.....	24
クラスの定義.....	24
クラス・インスタンスの作成.....	25
クラス・インスタンスへの属性の追加.....	25
クラス属性およびメソッドの定義.....	25
非表示変数.....	26
継承.....	26
第 3 章 IBM SPSS Modeler でのスクリプト.....	27
スクリプトの種類.....	27
ストリーム、スーパーノード・ストリーム、およびダイアグラム.....	27
ストリーム.....	27
スーパーノード・ストリーム.....	27
ダイアグラム.....	27
ストリームの実行.....	27
スクリプト・コンテキスト.....	28
既存のノードの参照.....	29
ノードの検索.....	29

プロパティを設定する.....	30
ノードの作成とストリームの変更.....	31
ノードの作成.....	31
ノードのリンクとリンク解除.....	32
ノードのインポート、置換、および削除.....	33
ストリーム内のノードのトラバース.....	34
項目の消去または削除.....	35
ノードに関する情報の入手.....	35
第 4 章スクリプト API.....	39
スクリプト API の概要.....	39
例 1: カスタム・フィルターを使用したノードの検索.....	39
例 2: ユーザーの権限に基づき、ディレクトリーまたはファイルの情報をユーザーが取得できるよ うにする.....	39
メタデータ: データに関する情報.....	40
生成されたオブジェクトへのアクセス.....	42
エラーの処理.....	44
ストリーム、セッション、およびスーパーノード・パラメーター.....	44
グローバル値.....	48
複数のストリームの処理: スタンドアロン スクリプト.....	49
第 5 章スクリプトのヒント.....	51
ストリーム実行の変更.....	51
ノードのループ.....	51
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・リポジトリー 内のオブジェクト へのアクセス.....	51
暗号化パスワードの生成.....	54
スクリプトの検査.....	54
コマンド・ラインからのスクリプト.....	55
旧リリースとの互換性.....	55
ストリーム実行結果へのアクセス.....	55
テーブル コンテンツ モデル.....	56
XML コンテンツ モデル.....	58
JSON コンテンツ モデル.....	59
列統計コンテンツ モデルおよびペアごとの統計コンテンツ モデル.....	61
第 6 章コマンド・ライン引数.....	65
ソフトウェアの起動.....	65
コマンド・ライン引数の使用.....	65
システムの引数.....	66
パラメーターの引数.....	67
サーバー接続の引数.....	68
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・リポジトリー 接続の引数.....	69
IBMSPPS Analytic Server 接続の引数.....	70
複数の引数の組み合わせ.....	70
第 7 章プロパティ・リファレンス.....	73
プロパティ・リファレンスの概要.....	73
プロパティのシンタックス.....	73
ノードおよびストリームのプロパティの例.....	75
ノードのプロパティの概要.....	75
共通のノード・プロパティ.....	76
第 8 章 Stream プロパティ.....	77
第 9 章入力ノードのプロパティ.....	81

入力ノードの共通プロパティ	81
asimport プロパティ	88
cognosimport ノードのプロパティ	89
databasenode プロパティ	93
datacollectionimportnode プロパティ	94
excelimportnode プロパティ	98
extensionimportnode プロパティ	99
fixedfilenode プロパティ	101
gsdata_import ノードのプロパティ	106
jsonimportnode のプロパティ	106
sasimportnode プロパティ	107
simgennode プロパティ	107
statisticsimportnode プロパティ	109
tm1odataimport ノードのプロパティ	109
tm1import ノードのプロパティ (廃止)	110
twcimport ノードのプロパティ	111
userinputnode プロパティ	112
variablefilenode プロパティ	113
xmlimportnode プロパティ	118
第 10 章レコード設定ノードのプロパティ	121
appendnode プロパティ	121
agggregatenode プロパティ	121
balancenode プロパティ	122
cplexoptnode プロパティ	123
derive_stbnode プロパティ	126
distinctnode プロパティ	128
extensionprocessnode プロパティ	129
mergenode プロパティ	131
rfmaggregatenode プロパティ	133
samplenode プロパティ	135
selectnode プロパティ	137
sortnode プロパティ	138
spacetimeboxes プロパティ	138
streamingtimeseries プロパティ	140
第 11 章フィールド設定ノードのプロパティ	151
anonymizenode プロパティ	151
autodatapreprenode プロパティ	152
astimeintervalsnode プロパティ	156
binningnode プロパティ	157
derivenode プロパティ	160
ensemblenode プロパティ	163
fillernode プロパティ	164
filternode プロパティ	165
historynode プロパティ	166
partitionnode プロパティ	167
reclassifynode プロパティ	168
reordernode プロパティ	169
reprojectnode プロパティ	170
restructurenode プロパティ	170
rfmanalysisnode プロパティ	171
settoflagnode プロパティ	172
statisticstransformnode プロパティ	173
timeintervalsnode プロパティ (廃止)	174
transposenode プロパティ	180
typenode プロパティ	182

第 12 章グラフ作成ノードのプロパティ	191
グラフ作成ノードの共通プロパティ	191
collectionnode プロパティ	192
distributionnode プロパティ	193
evaluationnode プロパティ	194
graphboardnode プロパティ	196
histogramnode プロパティ	201
mapvisualization プロパティ	202
multiplotnode プロパティ	207
plotnode プロパティ	208
timeplotnode プロパティ	211
eplotnode プロパティ	212
tsnnode プロパティ	213
webnode プロパティ	215
第 13 章モデル作成ノードのプロパティ	217
一般的なモデル作成ノードのプロパティ	217
anomalydetectionnode プロパティ	218
apriorinode プロパティ	219
associationrulesnode プロパティ	221
autoclassifiernode プロパティ	224
アルゴリズム・プロパティの設定	226
autoclusternode プロパティ	226
autonumericnode プロパティ	228
bayesnetnode プロパティ	230
c50node プロパティ	232
carmanode プロパティ	234
cartnode プロパティ	235
chaidnode プロパティ	238
coxregnode プロパティ	240
decisionlistnode プロパティ	243
discriminantnode プロパティ	244
extensionmodelnode プロパティ	246
factornode プロパティ	249
featureselectionnode プロパティ	251
genlinnode プロパティ	253
glmnode プロパティ	258
gle プロパティ	263
kmeansnode プロパティ	270
kmeansasnode プロパティ	271
knnnode プロパティ	272
kohonennode プロパティ	274
linearnode プロパティ	276
linearasnode プロパティ	277
logregnode プロパティ	279
lsvmnode プロパティ	284
neuralnetnode プロパティ	285
neuralnetworknode プロパティ	288
questnode プロパティ	290
randomtrees プロパティ	293
regressionnode プロパティ	295
sequencenode プロパティ	297
slrmnode プロパティ	299
statisticsmodelnode プロパティ	300
stpnode プロパティ	300
svmnode プロパティ	307

tcmnode プロパティ	308
ts プロパティ	313
treeas プロパティ	324
twostepnode プロパティ	326
twostepAS のプロパティ	327

第 14 章モデル・ナゲット・ノードのプロパティ..... 331

applyanomalydetectionnode プロパティ	331
applyapriorinode プロパティ	331
applyassociationrulesnode プロパティ	332
applyautoclassifiernode プロパティ	332
applyautoclusternode プロパティ	334
applyautonumericnode プロパティ	334
applybayesnetnode プロパティ	334
applyc50node プロパティ	335
applycarmanode プロパティ	335
applycartnode プロパティ	335
applychaidnode プロパティ	336
applycoxregnode プロパティ	336
applydecisionlistnode プロパティ	337
applydiscriminantnode プロパティ	337
applyextension プロパティ	338
applyfactornode プロパティ	339
applyfeatureselectionnode プロパティ	339
applygeneralizedlinearnode プロパティ	340
applyglmnode プロパティ	340
applygle プロパティ	341
applygmm のプロパティ	341
applykmeansnode プロパティ	342
applyknnnode プロパティ	342
applykohonenode プロパティ	342
applylinearnode プロパティ	342
applylinearasnode プロパティ	343
applylogregnode プロパティ	343
applylsvmnode プロパティ	343
applyneuralnetnode プロパティ	344
applyneuralnetworknode プロパティ	344
applyocsvmnode のプロパティ	345
applyquestnode プロパティ	345
applyrandomtrees プロパティ	346
applyregressionnode プロパティ	347
applyselflearningnode プロパティ	347
applysequencenode プロパティ	347
applysvmnode プロパティ	347
applystpnode プロパティ	348
applytcmnode プロパティ	348
applyts プロパティ	348
applytimeseriesnode プロパティ (廃止)	349
applytreeas プロパティ	349
applytwostepnode プロパティ	349
applytwostepAS のプロパティ	350
applyxgboosttreenode のプロパティ	350
applyxgboostlinearnode のプロパティ	350
hdbscannugget のプロパティ	350
kdeapply のプロパティ	351

第 15 章データベース・モデル作成ノードのプロパティ..... 353

Microsoft モデル作成ノードのプロパティ	353
Microsoft モデル作成ノードのプロパティ	353
Microsoft モデル・ナゲットのプロパティ	355
Oracle モデル作成ノードのプロパティ	357
Oracle モデル作成ノードのプロパティ	358
Oracle モデル・ナゲットのプロパティ	364
IBM NetezzaAnalytics モデル作成ノードのプロパティ	365
Netezza モデル作成ノードのプロパティ	365
Netezza モデル・ナゲットのプロパティ	380
第 16 章出力ノードのプロパティ	381
analysisnode プロパティ	381
dataauditnode プロパティ	382
extensionoutputnode プロパティ	384
kdeexport プロパティ	385
matrixnode プロパティ	386
meansnode プロパティ	389
reportnode プロパティ	391
setglobalsnode プロパティ	392
simevalnode プロパティ	393
simfitnode プロパティ	394
statisticsnode プロパティ	395
statisticsoutputnode プロパティ	396
tablenode プロパティ	396
transformnode プロパティ	399
第 17 章エクスポート・ノードのプロパティ	403
共通のエクスポート・ノード・プロパティ	403
asexport プロパティ	403
cognosexportnode プロパティ	404
databaseexportnode プロパティ	406
datacollectionexportnode プロパティ	411
excelexportnode プロパティ	412
extensionexportnode プロパティ	413
jsonexportnode のプロパティ	414
outputfilenode プロパティ	414
sasexportnode プロパティ	416
statisticsexportnode プロパティ	416
tm1odataexport ノードのプロパティ	416
tm1export ノードのプロパティ (廃止)	419
xmlexportnode プロパティ	421
第 18 章 IBM SPSS 統計 ノードのプロパティ	423
statisticsimportnode プロパティ	423
statisticstransformnode プロパティ	423
statisticsmodelnode プロパティ	424
statisticsoutputnode プロパティ	425
statisticsexportnode プロパティ	425
第 19 章 Python ノードのプロパティ	427
gmm のプロパティ	427
hdbscannode のプロパティ	428
kdemodel のプロパティ	429
kdeexport プロパティ	431
gmm のプロパティ	432
ocsvmnode のプロパティ	433
rfnode プロパティ	435

smotenode のプロパティ	437
tsnenode プロパティ	438
xgboostlinearnode のプロパティ	440
xgboosttreenode のプロパティ	441
第 20 章 Spark ノードのプロパティ	445
isotonicasnode プロパティ	445
kmeansasnode プロパティ	445
multilayerperceptronnode のプロパティ	446
xgboostasnode プロパティ	447
第 21 章 スーパーノードのプロパティ	451
付録 A ノード名のリファレンス	453
モデル・ナゲット名	453
重複するモデル名の回避	455
出力形式名	455
付録 B 従来のスクリプトから Python スクリプトへの移行	457
従来のスクリプトの移行の概要	457
一般的な差異	457
スクリプト・コンテキスト	457
コマンドと関数	457
リテラルとコメント	458
演算子	458
条件とループ	459
変数	460
ノード、出力、およびモデルの各タイプ	460
プロパティ名	461
ノードの参照	461
プロパティの取得と設定	461
ストリームの編集	462
ノード操作	463
ループ	463
ストリームの実行	464
ファイル・システムおよびリポジトリによるオブジェクトへのアクセス	465
ストリーム操作	465
モデルの操作	466
ドキュメント出力操作	466
従来のスクリプトと Python スクリプトのその他の違い	467
特記事項	469
商標	470
製品資料に関するご使用条件	470
索引	473

第1章 スクリプトとスクリプト言語

スクリプトの概要

IBM スポス モデラー のスクリプトは、ユーザー・インターフェースのプロセスを自動化する強力なツールです。スクリプトで、マウスやキーボードを使用した場合と同じ種類のアクションを実行できます。また、頻繁に繰り返したり手動で実行するのに時間がかかるタスクを自動化するために使用できます。

次の処理にスクリプトを使用できます。

- ストリームでノードを実行する特定の順序を指定する。
- CLEM (Control Language for Expression Manipulation) のサブセットを使用して、ノードにプロパティを設定したり、フィールドを作成したりする。
- 通常はユーザーとの対話によって実行される一連の操作 (例えば、モデルを作成してテストするなど) を自動化する。
- 十分なユーザーとの対話が必要な複雑な処理 (例えば、モデルの生成とテストを繰り返す交差検証手順など) を設定する。
- ストリームを操作する処理 (例えば、モデル学習ストリームの取得や実行、対応するモデル・テスト・ストリームの自動生成など) を設定する。

この章では、ストリームレベルのスクリプト、スタンドアロン スクリプト、および IBM スポス モデラー インターフェースのスーパーノード内のスクリプトに関する高度な説明と例を記述しています。スクリプト言語、シンタックス、およびコマンドは、以後の章で説明します。

注:

IBM スポス モデラー 内の IBM SPSS 統計 で作成されたスクリプトはインポートおよび実行できません。

スクリプトの種類

IBM スポス モデラー では、次の3種類のスクリプトが使用されます。

- **ストリーム・スクリプト** は、ストリーム・プロパティとして格納されるため、特定のストリームと一緒に保存およびロードされます。例えば、モデル・ナゲットの学習と適用のプロセスを自動化するストリーム・スクリプトを書くことができます。また、特定のストリームが実行されたときは常に、そのストリームのキャンパスの内容ではなく、スクリプトが実行されるように指定することもできます。
- **スタンドアロン スクリプト** は、どのストリームとも関連付けがなく、外部のテキスト・ファイルに保存されます。スタンドアロン スクリプトは、例えば、複数のストリームを一緒に操作する場合に使用できます。
- **スーパーノード スクリプト** は、スーパーノード ストリーム・プロパティとして格納されます。スーパーノード・スクリプトは、ターミナル・スーパーノードでのみ使用可能です。スーパーノード スクリプトは、スーパーノードの内容のシーケンスの実行を制御するのに使用できます。ターミナル以外の (入力またはプロセス) スーパーノードの場合、ストリーム・スクリプト内で直接、スーパーノードまたはスーパーノード内のノードにプロパティを定義できます。

ストリーム・スクリプト

スクリプトを使用して特定のストリーム内の操作をカスタマイズできます。また、スクリプトをそのストリームとともに保存することができます。ストリーム・スクリプトは、ストリーム内のターミナル・ノードの、特定の実行順序を指示するために使用されます。ストリーム・スクリプト ダイアログ・ボックスを使用して、現在のストリームとともに保存されているスクリプトを編集します。

「ストリームのプロパティ」ダイアログ・ボックスの「ストリーム・スクリプト」タブにアクセスするには

1. 「ツール」メニューから次の各項目を選択します。

ストリームのプロパティ>実行

2.「実行」タブをクリックして、現在のストリームのスクリプトの処理を行います。

ストリーム・スクリプトのダイアログ・ボックスの一番上にあるツールバー・アイコンを使用すると、次のような操作を実行できます。

- ウィンドウに既存のスタンドアロン スクリプトの内容をインポートする。
- スクリプトをテキスト・ファイルとして保存する。
- スクリプトを印刷する。
- デフォルト スクリプトを追加する。
- スクリプトを編集する (元に戻す、切り取り、コピー、貼り付けなど、標準的な編集機能を使用)。
- 現在のスクリプト全体を実行する。
- スクリプトから選択した行を実行する。
- 実行時にスクリプトを停止する (このアイコンは、スクリプトの実行時のみ使用可能になります)。
- スクリプトのシンタックスをチェックして、エラーが見つければ、ダイアログ・ボックスの下部ペインにそれを表示する。

注: バージョン 16.0 以降、スpos モデラー は Python スクリプト言語を使用します。16.0 よりも前のすべてのバージョンでは、スpos モデラー に固有のスクリプト言語 (現在では「従来のスクリプト」と呼ばれる) を使用していました。処理しているスクリプトのタイプに応じて、「実行」タブで「デフォルト (オプション スクリプト)」実行モードを選択してから、「Python」または「従来のもの」のいずれかを選択します。

ストリームの実行時にスクリプトを実行するかどうかを指定できます。ストリームの実行時に常に、スクリプトに指定された実行順序でこのスクリプトを実行するには、「このスクリプトを実行」を選択します。この設定により、ストリーム・レベルでの自動化を実現でき、素早いモデル構築が可能になります。ただし、デフォルトでは、ストリーム実行時にこのスクリプトは無視されます。「このスクリプトを無視」オプションを選択した場合でも、常にこのダイアログ・ボックスで直接スクリプトを実行できます。

スクリプト・エディターには、スクリプト・オーサリングを支援する以下の機能が用意されています。

- シンタックスの強調表示。キーワード、リテラル値 (文字列や数値など)、コメントが強調表示されます。
- 行番号付け。
- ブロックの一致。カーソルがプログラム・ブロックの開始位置に置かれると、対応する終了ブロックも強調表示されます。
- 自動入力候補表示。

シンタックスの強調表示で使用する色とテキストのスタイルは、IBM スpos モデラー の表示設定を使用してカスタマイズすることができます。この表示設定にアクセスするには、「ツール」>「オプション」>「ユーザー オプション」を選択して、「シンタックス」タブを選択します。

候補として表示されるシンタックス入力リストにアクセスするには、コンテキスト・メニューから「自動候補提示機能」を選択するか、Ctrl + スペースを押します。カーソル・キーを使用してリストを上下に移動し、Enter キーを押して、選択したテキストを挿入します。既存のテキストを変更せずに自動候補提示機能モードを終了するには、Esc キーを押します。

「デバッグ」タブには、デバッグ・メッセージが表示されます。このタブを使用すると、スクリプトの実行中にスクリプトの状態を評価することができます。「デバッグ」タブは、読み取り専用テキスト領域と 1 行の入力テキスト・フィールドから構成されています。テキスト領域には、スクリプトによって (例えばエラー・メッセージ・テキストを通じて) 標準出力または標準エラーのいずれかに送信されたテキストが表示されます。入力テキスト・フィールドは、ユーザーから入力を受け取ります。次に、この入力は、ダイアログ内で最後に実行されたスクリプトのコンテキスト (スクリプト・コンテキスト) 内で評価されます。テキスト領域には、ユーザーがコマンドのトレースを確認できるように、コマンドと結果の出力が表示されます。テキスト入力フィールドには、コマンド・プロンプト (従来のスクリプトの -->) が常に表示されます。

以下の場合、新しいスクリプト・コンテキストが作成されます。

- スクリプトを実行するには、「このスクリプトを実行」または「選択した行のみ実行」を使用します。

- スクリプト言語を変更した場合。

新しいスクリプト・コンテキストが作成されると、テキスト領域がクリアされます。

注: スクリプト ペインの外部でストリームを実行しても、スクリプト ペインのスクリプト・コンテキストは変更されません。この実行の一部として作成された変数の値は、スクリプト・ダイアログ・ボックス内には表示されません。

ストリーム・スクリプトの例:ニューラル・ネットワークの学習

ストリームは実行時に、ニューラル・ネットワーク・モデルの学習に使用できます。通常、モデルをテストするには、モデル作成ノードを実行してモデルをストリームに追加し、適切な接続を確立して、精度分析ノードを実行します。

IBM スポス モデラー スクリプトを使用すると、モデル・ナゲット作成後のテスト プロセスを自動化できます。例えば、デモ・ストリーム `druglearn.str` (IBM スポス モデラー インストール済み環境の `/Demos/streams/` フォルダーにあります) をテストするための以下のストリーム・スクリプトは、「ストリームのプロパティ」ダイアログ (「ツール」 > 「ストリーム・プロパティ」 > 「スクリプト」) から実行できます。

```
stream = modeler.script.stream()
neuralnetnode = stream.findByType("neuralnetwork", None)
results = []
neuralnetnode.run(results)
appliernode = stream.createModelApplierAt(results[0], "Drug", 594, 187)
analysisnode = stream.createAt("analysis", "Drug", 688, 187)
typenode = stream.findByType("type", None)
stream.linkBetween(appliernode, typenode, analysisnode)
analysisnode.run([])
```

このスクリプト例の各行について、次に説明します。

- 1 行目では、現在のストリームを指し示す変数を定義します。
- 2 行目では、スクリプトはニューラル・ネットワーク・ビルダー・ノードを検出します。
- 3 行目では、実行結果を格納できるリストを作成します。
- 4 行目では、ニューラル ネットワーク モデル ナゲットが作成されます。これは、3 行目で定義したリストに格納されます。
- 5 行目では、モデル ナゲットのモデル適用ノードが作成され、ストリーム領域に配置されます。
- 6 行目では、`Drug` という名前の精度分析ノードが作成されます。
- 7 行目では、スクリプトはデータ型ノードを検索します。
- 8 行目では、5 行目で作成したモデル適用ノードを、データ型ノードと分析ノードとの間で接続します。
- 最後に、精度分析ノードが実行されて、分析レポートが生成されます。

空の領域から、ストリームを初めから作成して実行するスクリプトを使用することも可能です。

Jython コードのサイズ制限

Jython は各スクリプトを Java バイトコードにコンパイルし、その Java バイトコードは Java 仮想マシン (JVM) によって実行されます。ただし、Java では単一のバイトコード・ファイルのサイズに制限があります。そのため、Jython がバイトコードをロードしようとした際に、JVM が異常終了することがあります。IBM スポス モデラー は、その発生を防ぐことができません。

Jython スクリプトの記述には、適正なコーディング手法 (変数や関数を使用して共通の中間値を計算することにより重複コードを最小にするなど) を使用するようになしてください。必要に応じて、コードを複数のソース・ファイルに分割するか、モジュールを使用してコードを定義して、これらが別々のバイトコード・ファイルにコンパイルされるようにしなければならない場合があります。

スタンドアロン スクリプト

「スタンドアロン スクリプト」ダイアログ・ボックスでは、テキスト・ファイルとして保存されるスクリプトを作成したり編集したりします。このダイアログ・ボックスには、ファイル名が表示されます。スクリプトのロード、保存、インポート、および実行の機能が備わっています。

スタンドアロン スクリプトのダイアログ・ボックスにアクセスするには

メイン・メニューから次の各項目を選択します。

ツール > スタンドアロンスクリプト

スタンドアロン スクリプトでは、ストリーム・スクリプトと同じツールバーやスクリプト・シンタックス 検査オプションを使用することができます。詳しくは、トピック [1 ページ](#)の『ストリーム・スクリプト』を参照してください。

スタンドアロン スクリプトの例 :モデルの保存とロード

スタンドアロン スクリプトは、ストリームを操作するときに役立ちます。2 種類のストリームがある場合を想定します。1 つはモデルを作成するストリームであり、もう 1 つはグラフを使用して最初のストリームと既存のデータ・フィールドから生成されたルール・セットを探索するストリームです。この場合のスタンドアロン スクリプトは次のようになります。

```
taskrunner = modeler.script.session().getTaskRunner()

# Modify this to the correct Modeler installation Demos folder.
# Note use of forward slash and trailing slash.
installation = "C:/Program Files/IBM/SPSS/Modeler/19/Demos/"

# First load the model builder stream from file and build a model
druglearn_stream = taskrunner.openStreamFromFile(installation + "streams/druglearn.str", True)
results = []
druglearn_stream.findByType("c50", None).run(results)

# Save the model to file
taskrunner.saveModelToFile(results[0], "rule.gm")

# Now load the plot stream, read the model from file and insert it into the stream
drugplot_stream = taskrunner.openStreamFromFile(installation + "streams/drugplot.str", True)
model = taskrunner.openModelFromFile("rule.gm", True)
modelapplier = drugplot_stream.createModelApplier(model, "Drug")

# Now find the plot node, disconnect it and connect the
# model applier node between the derive node and the plot node
derivencode = drugplot_stream.findByType("derive", None)
plotnode = drugplot_stream.findByType("plot", None)
drugplot_stream.disconnect(plotnode)
modelapplier.setPositionBetween(derivencode, plotnode)
drugplot_stream.linkBetween(modelapplier, derivencode, plotnode)
plotnode.setPropertyValue("color_field", "$C-Drug")
plotnode.run([])
```

注: スクリプト言語一般については、[15 ページ](#)の『スクリプト言語の概要』を参照してください。

スタンドアロン スクリプトの例 :特徴量選択モデルの生成

この例では、空の領域から特徴量選択モデルを生成するストリームを構築し、そのモデルに適用して、指定された対象に関連するもっとも重要な上位 15 のフィールドを表示するテーブルを作成します。

```
stream = modeler.script.session().createProcessorStream("featureselection",
True)

statisticsimportnode = stream.createAt("statisticsimport", "Statistics
File", 150, 97)
statisticsimportnode.setPropertyValue("full_filename", "$CLEO_DEMOS/
customer_dbase.sav")

typenode = stream.createAt("type", "Type", 258, 97)
typenode.setKeyedPropertyValue("direction", "response_01", "Target")
```

```

featureselectionnode = stream.createAt("featureselection", "Feature
Selection", 366, 97)
featureselectionnode.setPropertyValue("top_n", 15)
featureselectionnode.setPropertyValue("max_missing_values", 80.0)
featureselectionnode.setPropertyValue("selection_mode", "TopN")
featureselectionnode.setPropertyValue("important_label", "Check Me Out!")
featureselectionnode.setPropertyValue("criteria", "Likelihood")

stream.link(statisticsimportnode, typenode)
stream.link(typenode, featureselectionnode)
models = []
featureselectionnode.run(models)

# Assumes the stream automatically places model apply nodes in the stream
applynode = stream.findByType("applyfeatureselection", None)
tablenode = stream.createAt("table", "Table", applynode.getXPosition() + 96,
applynode.getYPosition())
stream.link(applynode, tablenode)
tablenode.run([])

```

このスクリプトで、データを読み込む入力ノードを作成し、**response_01** フィールドの役割を **Target** に設定するデータ型ノードを使用し、その後特徴量選択ノードを作成して実行します。また、読みやすいレイアウトになるように、ストリーム領域で各ノードを接続し、配置します。その後、作成されるモデル・ナゲットがテーブル・ノードへ接続されます。テーブル・ノードでは、**selection_mode** プロパティと **top_n** プロパティに設定されたとおりに、もっとも重要な上位 15 フィールドが一覧表示されます。詳しくは、トピック [251 ページの『featureselectionnode プロパティ』](#) を参照してください。

スーパーノード・スクリプト

IBM スポス モデラー のスクリプト言語を使用して、スクリプトを作成し、任意のターミナル・スーパーノード内に保存できます。これらのスクリプトはターミナル・スーパーノードにのみ使用でき、テンプレートストリームの作成時、およびスーパーノードの内容に特定の実行順序を指定する際に使用できます。スーパーノード・スクリプトを使用すると、ストリーム内で複数のスクリプトを実行することもできます。

例えば、複雑なストリームで実行の順序を指定する必要があり、スーパーノードには、散布図ノードで使われる新しいフィールドを作成する前に実行される必要のあるグローバルの設定ノードを含む、いくつかのノードがあるとします。この場合、まずグローバルの設定ノードを実行するスーパーノード・スクリプトを作成できます。このノードが計算する平均や標準偏差などの値は、散布図ノードを実行するときに使用します。

スーパーノード・スクリプト内では、ほかのスクリプトの場合と同様の方法で、ノード・プロパティを指定できます。また、ストリーム・スクリプトから直接に、任意のスーパーノードまたはカプセル化されたノードのプロパティを変更または定義することもできます。詳しくは、トピック [451 ページの『第 21 章 スーパーノードのプロパティ』](#) を参照してください。この手法は、ソース スーパーノード、プロセス スーパーノード、およびターミナル・スーパーノードに適用できます。

注: 独自のスクリプトを実行することができるのはターミナル・スーパーノードの場合だけなので、「スーパーノード」ダイアログ・ボックスの「スクリプト」タブは、ターミナル・スーパーノードの場合にだけ利用可能です。

メインキャンバスから「スーパーノード・スクリプト」ダイアログ・ボックスを開くには

ストリームキャンバスでターミナル・スーパーノードを選択して、「スーパーノード」メニューから次の項目を選択します。

スーパーノード・スクリプト...

ズーム・インしたスーパーノード・キャンバスから「スーパーノード・スクリプト」ダイアログ・ボックスを開くには

スーパーノード領域を右クリックして表示されるコンテキスト・メニューから、次の項目を選択します。
スーパーノード・スクリプト...

スーパーノード・スクリプトの例

次のスーパーノード・スクリプトでは、スーパーノード内のターミナル・ノードが実行されるべき順序が宣言されます。この順序によって、まずグローバルの設定ノードが実行されて、別のノードを実行したときに、このノードによって算出される値が使用されるようになります。

```
execute 'Set Globals'  
execute 'gains'  
execute 'profit'  
execute 'age v. $CC-pep'  
execute 'Table'
```

スーパーノードのロックとロック解除

次の例は、スーパーノードのロックおよびロック解除の方法を説明しています。

```
stream = modeler.script.stream()  
superNode=stream.findByID('id854RNTSD5MB')  
# unlock one super node  
print 'unlock the super node with password abcd'  
if superNode.unlock('abcd'):  
    print 'unlocked.'  
else:  
    print 'invalid password.'  
# lock one super node  
print 'lock the super node with password abcd'  
superNode.lock('abcd')
```

ストリームでのループと条件付き実行

バージョン 16.0 以降、SPSS モデラー では、スクリプト言語で直接指示を作成するのではなく、さまざまなダイアログ・ボックスで値を選択することによって、ストリーム内でいくつかの基本的なスクリプトを作成することができます。この方法で作成できるスクリプトの 2 つの主なタイプは、単純ループと、条件が満たされた場合にノードを実行する方法です。

ストリーム内でループ規則と条件付き実行規則の両方を組み合わせることができます。例えば、世界中の製造業者の自動車の販売に関連したデータがあるとして、ストリーム内のデータを処理するループを設定して、製造業者の国別に詳細を識別し、モデル別の販売数、製造業者別およびエンジン・サイズ別の排気ガスレベルなどの詳細を示すさまざまなグラフにデータを出力することができます。ヨーロッパの情報のみに関心がある場合は、アメリカとアジアの製造業者用のグラフが作成されないようにする条件をループに追加することもできます。

注：ループと条件付き実行は、いずれもバックグラウンドのスクリプトに基づいているため、これらはストリームの実行時にストリーム全体にのみ適用されます。

- **ループ** ループを使用して、反復タスクを自動化できます。例えば、指定の数のノードをストリームに追加し、追加するたびに 1 つのノード・パラメーターを変更することができます。あるいは、以下の例のように、ストリームまたは枝を指定の数だけ繰り返し実行することを制御できます。
 - ストリームを指定の回数実行し、実行するたびにソースを変更する。
 - ストリームを指定の回数実行し、実行するたびに変数の値を変更する。
 - ストリームを指定の回数実行し、実行するたびに 1 つの追加フィールドを入力する。
 - モデルを指定の回数構築し、毎回モデル設定を変更する。
- **条件付き実行** これを使用すると、事前定義した条件に基づいて、どのようにターミナル・ノードを実行するかを制御できます。例えば、以下のようにします。

- 指定の値が true か false かに基づいて、ノードを実行するかどうかを制御する。
- ノードのループを並行して実行するのか、順次に実行するのかを定義する。

ループと条件付き実行は両方とも、「ストリームのプロパティ」ダイアログ・ボックス内の「実行」タブで設定します。条件またはループ要件に使用されているノードは、追加の記号が付加された状態でストリーム・キャンバスに表示され、ループおよび条件付き実行に関与していることが示されます。

「実行」タブには、以下の3つのいずれかの方法でアクセスできます。

- メイン・ダイアログ・ボックスの上部にあるメニューを使用する。
 1. 「ツール」メニューから次の各項目を選択します。
 - 「ストリームのプロパティ」 > 「実行」
 2. 「実行」タブをクリックして、現在のストリーム用のスクリプトを処理します。
- ストリーム内から。
 1. ノードを右クリックして「**ループ/条件付き実行**」を選択する。
 2. 該当するサブメニューのオプションを選択します。
- メイン・ダイアログ・ボックスの上部にあるグラフィック・ツールバーで、ストリーム・プロパティのアイコンをクリックする。

ループまたは条件付き実行の詳細を初めて設定する場合は、「実行」タブで「**ループ/条件付き実行**」実行モードを選択してから、「**条件**」または「**ループ**」サブタブを選択します。

ストリームでのループ

ループを使用すると、ストリーム内の反復タスクを自動化できます。例えば、以下のようになります。

- ストリームを指定の回数実行し、実行するたびにソースを変更する。
- ストリームを指定の回数実行し、実行するたびに変数の値を変更する。
- ストリームを指定の回数実行し、実行するたびに1つの追加フィールドを入力する。
- モデルを指定の回数構築し、毎回モデル設定を変更する。

満たすべき条件は、ストリームの「実行」タブの「**ループ**」サブタブで設定します。サブタブを表示するには、「**ループ/条件付き実行**」実行モードを選択します。

定義するループ要件は、「**ループ/条件付き実行**」実行モードが設定されている場合は、ストリーム実行時に有効になります。オプションで、ループ要件のスクリプト・コードを生成し、**貼り付け ...** をクリックしてスクリプト・エディターに貼り付けることができます。「**ループ**」サブタブの右下隅に表示されます。メインの「実行」タブが変更され、タブの上部にスクリプトがある **デフォルト (オプション・スクリプト)** 実行モードが表示されます。つまり、スクリプト・エディターで詳細にカスタマイズ可能なスクリプトを生成する前に、ダイアログ・ボックスのさまざまなループ・オプションを使用してループ構造を定義できます。**貼り付け ...** をクリックすると、定義した条件付き実行要件は、生成されたスクリプトにも表示されます。

重要: スポス モデラー ストリームで設定したループ変数は、このストリームを 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 ジョブで実行する場合、上書きされる可能性があります。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 ジョブ エディタのエントリは、スポス モデラー のエントリによって上書きされるからです。例えば、ループごとに異なる出力ファイルを作成するようにストリーム内のループ変数を設定した場合、これらのファイルは スポス モデラー 内で正しく命名されますが、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 Deployment Manager の「結果」タブに入力された固定エントリによって上書きされます。

ループをセットアップするには

1. ストリームで実行するメインのループ構造を定義する反復キーを作成します。詳しくは、[反復キーの作成](#)を参照してください。
2. 必要に応じて、1つ以上の反復変数を定義します。詳しくは、[反復変数の作成](#)を参照してください。

3. 作成した反復および変数が、サブタブの本文に表示されます。デフォルトで、反復は表示されている順に実行されます。反復をリスト内で上または下に移動するには、反復をクリックして選択し、サブタブの右側の列にある上矢印または下矢印を使用して順序を変更します。

ストリームでのループのための反復キーの作成

反復キーを使用して、ストリームで実行するメインのループ構造を定義します。例えば、自動車販売を分析している場合は、ストリーム・パラメーター製造国を作成し、これを反復キーとして使用できます。ストリームを実行すると、このキーは、反復のたびにデータ内でそれぞれ異なる国の値に設定されます。「反復キーの定義」ダイアログ・ボックスを使用して、キーを設定します。

ダイアログ・ボックスを開くには、**反復キー ...** を選択します。「ループ」サブタブの左下隅にあるボタンをクリックするか、ストリーム内の任意のノードを右クリックして、**ループ/条件付き実行 > 反復キーの定義 (フィールド)** または **ループ/条件付き実行 > 反復キーの定義 (値)** のいずれかを選択します。ストリームからダイアログ・ボックスを開くと、ノードの名前などの一部のフィールドが自動的に入力されます。

反復キーを設定するには、以下のフィールドに入力します。

反復対象。 次のいずれかのオプションを選択できます。

- **ストリーム パラメーター - フィールド。** このオプションは、既存のストリーム・パラメーターの値を、指定した各フィールドに順に設定するループを作成する場合に使用します。
- **ストリーム パラメーター - 値。** このオプションは、既存のストリーム・パラメーターの値を、指定した各値に順に設定するループを作成する場合に使用します。
- **ノード プロパティ - フィールド。** このオプションは、ノード・プロパティの値を、指定した各フィールドに順に設定するループを作成する場合に使用します。
- **ノード プロパティ - 値。** このオプションは、ノード・プロパティの値を、指定した各値に順に設定するループを作成する場合に使用します。

設定内容。 ループが実行されるたびに値が設定される項目を選択します。次のいずれかのオプションを選択できます。

- **パラメーター。** 「ストリーム パラメーター - フィールド」または「ストリーム パラメーター - 値」を選択した場合にのみ使用可能です。使用可能なリストから、必要なパラメーターを選択します。
- **ノード。** 「ノード プロパティ - フィールド」または「ノード プロパティ - 値」を選択した場合にのみ使用可能です。ループをセットアップするノードを選択します。参照ボタンをクリックして「ノード選択」ダイアログを開き、目的のノードを選択します。リストされているノードが多すぎる場合は、入力ノード、プロセス・ノード、グラフ作成ノード、モデル作成ノード、出力ノード、エクスポート・ノード、またはモデル・ノードの適用のいずれかのカテゴリ別にノードを表示するように、表示をフィルタリングできます。
- **プロパティ。** 「ノード プロパティ - フィールド」または「ノード プロパティ - 値」を選択した場合にのみ使用可能です。使用可能なリストからノードのプロパティを選択します。

使用するフィールド。 「ストリーム パラメーター - フィールド」または「ノード プロパティ - フィールド」を選択した場合にのみ使用可能です。反復値を提供するために使用するノード内のフィールドを選択します。次のいずれかのオプションを選択できます。

- **ノード。** 「ストリーム パラメーター - フィールド」を選択した場合にのみ使用可能です。ループを設定する詳細を含むノードを選択します。参照ボタンをクリックして「ノード選択」ダイアログを開き、目的のノードを選択します。リストされているノードが多すぎる場合は、入力ノード、プロセス・ノード、グラフ作成ノード、モデル作成ノード、出力ノード、エクスポート・ノード、またはモデル・ノードの適用のいずれかのカテゴリ別にノードを表示するように、表示をフィルタリングできます。
- **フィールド リスト。** 右側の列にあるリスト・ボタンをクリックして、「フィールドの選択」ダイアログ・ボックスを表示します。このダイアログ・ボックスで、反復データを提供するノード内のフィールドを選択します。詳しくは、[9 ページの『反復のためのフィールドの選択』](#) を参照してください。

使用する値。 「ストリーム パラメーター - 値」または「ノード プロパティ - 値」を選択した場合にのみ使用可能です。選択したフィールド内で、反復値として使用する値 (複数可) を選択します。次のいずれかのオプションを選択できます。

- **ノード**. 「**ストリーム パラメーター - 値**」を選択した場合にのみ使用可能です。ループを設定する詳細を含むノードを選択します。参照ボタンをクリックして「ノード選択」ダイアログを開き、目的のノードを選択します。リストされているノードが多すぎる場合は、入力ノード、プロセス・ノード、グラフ作成ノード、モデル作成ノード、出力ノード、エクスポート・ノード、またはモデル・ノードの適用のいずれかのカテゴリー別にノードを表示するように、表示をフィルタリングできます。
- **フィールド リスト**. 反復データを提供するためのノード内のフィールドを選択します。
- **値リスト**. 右側の列にあるリスト・ボタンをクリックして、「値の選択」ダイアログ・ボックスを表示します。このダイアログ・ボックスで、反復データを提供するフィールド内の値を選択します。

ストリームでのループのための反復変数の作成

反復変数を使用して、ループが実行されるたびに、ストリーム内の選択したノードのストリーム・パラメーターまたはプロパティの値を変更できます。例えば、ストリーム・ループが自動車販売データを分析していて、製造国を反復キーとして使用している場合に、モデル別の販売を示すグラフ出力と、排気ガス情報を示すグラフ出力があるとします。このような場合に、結果グラフごとに新しいタイトル(スウェーデンの自動車排気ガス や日本のモデル別自動車販売 など)を作成する反復変数を作成できます。「反復変数の定義」ダイアログ・ボックスを使用して、必要な変数を設定します。

ダイアログ・ボックスを開くには、**変数の追加 ...** を選択します。「ループ」サブタブの左下隅にあるボタンをクリックするか、ストリーム内の任意のノードを右クリックして**ループ/条件付き実行 > 反復変数の定義**を選択します。

反復変数を設定するには、以下のフィールドに入力します。

変更. 修正する属性の種類を選択します。「**ストリーム パラメーター**」または「**ノード プロパティ**」を選択します。

- 「**ストリーム パラメーター**」を選択した場合、必要なパラメーターを選択してから、以下のいずれかのオプションを使用して(ストリームで使用可能な場合)、ループを反復するたびにそのパラメーターに設定する値を定義します。
 - **グローバル変数**. ストリーム・パラメーターを設定するグローバル変数を選択します。
 - **テーブル出力セル**. ストリーム・パラメーターをテーブル出力セルの値に設定するには、リストからテーブルを選択して、使用する「**行**」と「**列**」を入力します。
 - **手動で入力**. このオプションは、このパラメーターが反復のたびに取る値を手動で入力する場合に選択します。「ループ」サブタブに戻ると、必要なテキストを入力する新しい列が作成されます。
- 「**ノード プロパティ**」を選択した場合は、必要なノードといずれかのプロパティを選択してから、そのプロパティに使用する値を設定します。以下のオプションの1つを使用して、新しいプロパティ値を設定します。
 - **単独**. プロパティ値は、反復キー値を使用します。詳しくは、[8 ページの『ストリームでのループのための反復キーの作成』](#)を参照してください。
 - **語幹の接頭辞として**. 「語幹」フィールドに入力する内容の接頭辞として反復キー値を使用します。
 - **語幹の接尾辞として**. 「語幹」フィールドに入力する内容の接尾辞として反復キー値を使用します。

接頭辞または接尾辞のオプションを選択した場合は、「語幹」フィールドに追加テキストを追加するように求められます。例えば、反復キー値が製造国 で、「**語幹の接頭辞として**」を選択した場合は、このフィールドに - モデル別の販売 と入力できます。

反復のためのフィールドの選択

反復を作成する場合は、「フィールドの選択」ダイアログ・ボックスを使用して、1つ以上のフィールドを選択できます。

ソート基準: 以下のいずれかのオプションを選択することにより、使用可能なフィールドを表示用にソートすることができます。

- **ファイル順**: データ ストリームから現在のノードに渡された順に各フィールドを表示します。
- **名前**: 各フィールドをアルファベット順にソートして表示します。

- **タイプ:** 各フィールドを測定の尺度順にソートして表示します。特定の尺度のフィールドを選択する場合に役立ちます。

リストからフィールドを1回に1つずつ選択するか、または Shift キーまたは Ctrl キーを押しながら複数のフィールドを選択します。また、リストの下ボタンを使用して、尺度に基づいて複数のフィールドを選択したり、テーブル中のすべてのフィールドを選択または選択解除することができます。

選択可能なフィールドは、使用しているストリーム・パラメーターまたはノード・プロパティに適切なフィールドのみが表示されるようにフィルタリングされていることに注意してください。例えば、ストレージ・タイプが文字列のストリーム・パラメーターを使用している場合は、ストレージ・タイプが文字列のフィールドのみが表示されます。

ストリームでの条件付き実行

条件付き実行では、定義するストリーム内容の一致条件に基づいて、ターミナル・ノードを実行する方法を制御できます。例えば、以下のようにできます。

- 指定の値が true か false かに基づいて、ノードを実行するかどうかを制御する。
- ノードのループを並行して実行するのか、順次に実行するのかを定義する。

満たすべき条件は、ストリームの「実行」タブの「条件」サブタブで設定します。サブタブを表示するには、「**ループ/条件付き実行**」実行モードを選択します。

定義する条件付き実行要件は、「**ループ/条件付き実行**」実行モードが設定されている場合は、ストリーム実行時に有効になります。オプションで、条件付き実行要件のスクリプト・コードを生成し、**貼り付け ...** をクリックしてスクリプト・エディターに貼り付けることができます。「条件付き」サブタブの右下隅に表示されます。メインの「実行」タブの表示が変更され、タブの上部にスクリプトが表示された **デフォルト (オプション・スクリプト)** 実行モードが表示されます。つまり、スクリプト・エディターで詳細にカスタマイズ可能なスクリプトを生成する前に、ダイアログ・ボックスのさまざまなループ・オプションを使用して条件を定義できます。**貼り付け ...** をクリックすると、定義したループ要件は、生成されたスクリプトにも表示されます。

条件をセットアップするには以下を行います。

1. 「条件」サブタブの右側の列で「新規条件の追加」ボタン  をクリックして、「**条件実行式の追加**」ダイアログボックスを開きます。このダイアログで、ノードを実行するために満たす必要がある条件を指定します。
2. 「**条件実行式の追加**」ダイアログボックスで、以下のオプションを指定します。
 - a. **ノード.** 条件付き実行を設定するノードを選択します。参照ボタンをクリックして「ノード選択」ダイアログを開き、目的のノードを選択します。リストされているノードが多すぎる場合は、エクスポート・ノード、グラフ作成ノード、モデル作成ノード、または出力ノードのいずれかのカテゴリ別にノードを表示するように、表示をフィルタリングできます。
 - b. **条件の基準.** ノードを実行するために満たす必要がある条件を指定します。「**ストリーム パラメーター**」、「**グローバル変数**」、「**テーブル出力セル**」、または「**常に True**」の4つのオプションのいずれか1つを選択できます。ダイアログ・ボックスの下半分に入力する詳細は、選択する条件によって異なります。
 - **ストリーム パラメーター.** 使用可能なリストからパラメーターを選択してから、そのパラメーターの「**演算子**」を選択します。例えば、演算子は「より大きい」、「等しい」、「より小さい」、「間」などです。次に、演算子に応じて「**値**」か、最小値および最大値を入力します。
 - **グローバル変数.** 使用可能なリストから変数を選択します。例えば、「平均」、「合計」、「最小値」、「最大値」、または「標準偏差」がリストに含まれている可能性があります。次に、「**演算子**」と必要な値を選択します。
 - **テーブル出力セル.** 使用可能なリストからテーブル・ノードを選択して、テーブルの「**行**」と「**列**」を選択します。次に、「**演算子**」と必要な値を選択します。
 - **常に True.** ノードを常に実行する必要がある場合は、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場合は、さらに選択するパラメーターはありません。

3. 必要なすべての条件を設定するまで、ステップ 1 と 2 を必要なだけ繰り返します。選択したノードと、ノードが実行される前に満たすべき条件が、サブタブの本体部分の「**実行ノード**」列と、「**値の設定条件 (真の場合に値を設定)**」列に表示されます。
4. デフォルトで、ノードと条件は表示されている順に実行されます。ノードと条件をリスト内で上または下に移動するには、ノードと条件をクリックして選択し、サブタブの右側の列にある上矢印または下矢印を使用して順序を変更します。

さらに、「条件」サブタブの下部にある以下のオプションを設定できます。

- **すべてを順番に評価。** このオプションは、各条件をサブタブに表示されている順序で評価する場合に選択します。条件が「True」のすべてのノードは、すべての条件が評価されてから 1 回だけ実行されます。
- **一度に 1 つずつ実行。** 「すべてを順番に評価」が選択されている場合にのみ使用できます。このオプションを選択すると、条件が「True」と評価される場合、その条件に関連付けられているノードは、次の条件が評価される前に実行されます。
- **最初のヒットまで評価。** このオプションを選択すると、指定した条件から「True」の評価が返される最初のノードのみが実行されます。

スクリプトの実行と中断

その他多くの方法でスクリプトを実行できます。例えば、ストリーム・スクリプトまたはスタンドアロンのスクリプトのダイアログで、「このスクリプトを実行」ボタンをクリックすると、完全なスクリプトを実行します。



図 1. 「このスクリプトを実行」ボタン

「選択した行」ボタンをクリックすると、スクリプト内で選択した 1 行または隣接する行のブロックを実行します。



図 2. 「選択した行を実行」ボタン

スクリプトの実行は、次のいずれかの方法で行います。

- ストリーム・スクリプトまたはスタンドアロン・スクリプトのダイアログ・ボックスの「このスクリプトを実行」または「選択した行を実行」をクリックします。
- デフォルトの実行方法として「このスクリプトを実行」が設定されているストリームを実行する。
- 起動時にインタラクティブ・モードで `-execute` フラグを使用します。詳しくは、トピック [65 ページ](#) の『[コマンド・ライン引数の使用](#)』を参照してください。

注：「スーパーノード」ダイアログ・ボックスで「このスクリプトを実行」を選択しているかぎり、スーパーノード・スクリプトは、スーパーノードの実行時に実行されます。

スクリプト実行の中断

「ストリーム・スクリプト」ダイアログ・ボックスのツールバーにある赤い中止ボタンは、スクリプト実行時に有効になります。このボタンを使用すると、スクリプトおよび現在のストリームの実行を中止することができます。

検索と置換

「検索/置換」ダイアログ・ボックスは、スクリプト・エディター、CLEM 式ビルダーなど、スクリプトまたは式のテキストを編集する場合、またはレポート・ノードでテンプレートを定義する場合に使用できます。これらの領域のいずれかでテキストを編集する場合、**Ctrl + F** キーを押してダイアログ・ボックスにアクセスし、カーソルがテキスト領域にフォーカスしていることを確認します。「置換」ノードを使用して

いる場合、例えば、「設定」タブのテキスト領域から、または CLEM 式ビルダーのテキスト・フィールドからダイアログ・ボックスにアクセスできます。

1. テキスト領域内にカーソルを置いて、Ctrl + F キーを押して「検索/置換」ダイアログ・ボックスにアクセスします。
2. 検索するテキストを入力するか、最近検索した項目のドロップダウン・リストから選択します。
3. 置換テキストがある場合は、入力します。
4. 「次を検索」をクリックして、検索を開始します。
5. 「置換」をクリックして現在の選択内容を置換するか、「すべてを置換」をクリックしてすべてまたは選択したインスタンスを更新します。
6. 各操作が終了すると、ダイアログ・ボックスが閉じます。テキスト領域で F3 を押すと最後の検索操作が繰り返され、または Ctrl + F キーを押すとダイアログに再度アクセスします。

検索オプション

大/小文字の一致。 検索操作で大/小文字を区別するかどうかを指定します。例えば、マイヴァルが *myVar* と一致するかどうかを指定します。この設定に関係なく、置換テキストは常に入力したとおりに挿入されます。

言葉だけだ 検索操作が単語内に埋め込まれたテキストと一致するかどうかを指定します。このオプションを選択すると、*spider* に関する検索は、*spiderman* または *spider-man* に一致しません。

正規表現。 正規表現の構文を使用するかどうかを指定します (次のセクションを参照)。このオプションを選択すると、「語全体のみ」オプションは無効化され、その値は無視されます。

選択されたテキストのみ: 「すべてを置換」オプションを使用する場合、検索の範囲を制御します。

正規表現シンタックス

正規表現を使用すると、タブまたは改行文字などの特殊文字、*a* から *d* までなど文字のクラスまたは範囲、行の開始または終了などの境界について検索することができます。次の種類の表現がサポートされています。

表 1. 文字の一致	
文字	一致
x	文字 x
\\	円記号
\\0n	8 進法の値を持つ文字 0n (0 <= n <= 7)
\\0nn	8 進法の値を持つ文字 0nn (0 <= n <= 7)
\\0mnn	8 進法の値を持つ文字 0mnn (0 <= m <= 3, 0 <= n <= 7)
¥xhh	16 進法の値を持つ文字 0xhh
¥uhhhh	16 進法の値を持つ文字 0xhhhh
\\t	タブ文字 ('¥u0009')
\\n	改行文字 ('¥u000A')
\\r	復帰文字 ('¥u000D')
\\f	改ページ文字 ('¥u000C')
\\a	アラート (ベル) 文字 ('¥u0007')
\\e	エスケープ文字 ('¥u001B')
¥cx	x に対応する制御文字

表 2. 文字クラスの一一致	
文字クラス	一致
[abc]	a、b、または c (単純クラス)
「^abc	a、b、または c 以外の文字 (減法)
「a-zA-Z]	a から z または A から Z の各文字 (範囲)
「a-d[m-p]]	a から d、または m から p (和集合)。または、「a-dm-p」と指定することもできます
「a-z&&[def]]	a から z、および d、e、または f (交差)
「a-z&&[^bc]]	a から z のうち、b と c を除いたもの (差集合)。または、「ad-z」と指定することもできます
「a-z&&[^m-p]]	a から z のうち、m から p までを除いたもの (差集合)。または、「a-lq-z」と指定することもできます

表 3. 事前設定された文字クラス	
事前設定された文字クラス	一致
.	任意の文字 (行末に一致する場合または一致しない場合があります)
\d	任意の数字: [0-9]
\D	数字以外: [^0-9]
\s	空白文字: [\t\n\x0B\f\r]
\S	空白文字以外: [^ \t\n\x0B\f\r]
\w	ワード文字: [a-zA-Z_0-9]
\W	非ワード文字: [^a-zA-Z_0-9]

表 4. 境界の一一致	
境界の一一致	一致
^	行頭
\$	行末
\b	語の境界
\B	語以外の境界
\A	入力の開始
\Z	最後の行末以外の入力の終了
\z	入力の終了

第2章 スクリプト言語

スクリプト言語の概要

IBM SPSS Modeler のスクリプト機能を使用すると、SPSS Modeler ユーザー・インターフェースで動作し、出力オブジェクトを操作し、コマンド・シンタックスを実行するスクリプトを作成できます。SPSS Modeler 内から直接スクリプトを実行できます。

IBM SPSS Modeler のスクリプトは、スクリプト言語 Python で作成されています。IBM SPSS Modeler で使用される Python の Java ベースの実装を Jython と呼びます。このスクリプト言語は、以下の機能で構成されています。

- ノード、ストリーム、プロジェクト、出力、およびその他の IBM SPSS Modeler オブジェクトを参照する形式
- 上記オブジェクトを操作するのに使用されるスクリプト ステートメントまたはコマンドのセット
- 変数、パラメーター、およびその他のオブジェクトに値を設定するためのスクリプト式の言語
- コメント、行の継続、およびリテラルテキストのブロックのサポート

以下のセクションでは、Python スクリプト言語、Python の Jython 実装、および IBM SPSS Modeler 内でスクリプトを使い始めるための基本シンタックスについて説明します。特定のプロパティとコマンドについての情報は、以後のセクションにあります。

Python と Jython

Jython は、Python スクリプト言語の実装の 1 つであり、Java 言語で記述され、Java プラットフォームと統合されています。Python は強力なオブジェクト指向スクリプト言語です。Jython は、成熟したスクリプト言語の生産性向上機能を備え、Python とは異なり、Java 仮想マシン (JVM) をサポートするすべての環境で動作します。そのため、プログラムの作成時に JVM の Java ライブラリーを使用することができます。Jython を使用すると、この違いを利用できると同時に、Python 言語の構文とほとんどの機能を使用できます。

スクリプト言語であるため、Python (およびその Jython 実装) は習得が容易で効率的にコーディングできるほか、動作するプログラムの作成に最小限の構造しか必要としません。コードは対話式で (一度に 1 行) 入力することができます。Python はインタプリター式のスクリプト言語であり、Java にあるプリコンパイルの段階がありません。Python プログラムは単なるテキスト・ファイルであり、(構文エラーがないかどうか構文解析された後に) 入力として解釈されます。単純な式 (定義済みの値など) のほか、複雑な操作 (関数定義など) もただちに実行され、使用可能になります。コードに対して行った変更を迅速にテストすることができます。しかし、スクリプトの解釈には不利な点もあります。例えば、未定義の変数を使用してもコンパイラー・エラーにならないため、その変数を使用するステートメントが実行される場合に限り、その実行のときに検出されます。この場合は、プログラムを編集して実行し、エラーをデバッグすることができます。

Python では、データやコードも含め、あらゆるものをオブジェクトとして扱います。したがって、それらのオブジェクトを一連のコードで操作することができます。一部の型 (数値や文字列など) はオブジェクトではなく値と見なすと便利ですが、この扱いは Python でもサポートされています。サポートされているヌル値が 1 つあります。このヌル値には予約名 None が割り当てられています。

Python スクリプトおよび Jython スクリプトの概要やスクリプト例については、<http://www.ibm.com/developerworks/java/tutorials/j-jython1/j-jython1.html> および <http://www.ibm.com/developerworks/java/tutorials/j-jython2/j-jython2.html> を参照してください。

Python スクリプト

Python スクリプト言語の以下のガイドでは、IBM SPSS Modeler でスクリプトを作成する場合に使用される可能性が高いコンポーネントの概要と、概念やプログラミングの基礎について取り上げます。これによ

り、IBM SPSS Modeler 内で使用する Python スクリプトの開発を始めるのに十分な知識を得ることができます。

操作

割り当ては等号 (=) を使用して行われます。例えば、値「3」を「x」という変数に割り当てるには、次のステートメントを使用します。

```
x = 3
```

等号は、文字列型のデータを変数に代入する場合にも使用されます。例えば、値「a string value」を「y」という変数に代入するには、以下のステートメントを使用します。

```
y = "a string value"
```

次の表に、よく使用される比較演算子および数値演算子と、その説明を示します。

表 5. 一般的な比較演算子および数値演算子	
操作	説明
x < y	x が y より小さいかどうか
x > y	x が y より大きいかどうか
x <= y	x が y 以下かどうか
x >= y	x が y 以上かどうか
x == y	x が y と等しいかどうか
x != y	x が y と等しくないかどうか
x <> y	x が y と等しくないかどうか
x + y	y を x に加算する
x - y	y を x から減算する
x * y	x に y を乗算する
x / y	x を y で除算する
x ** y	x を y 乗する

リスト

リストは、一連の要素です。リストには任意の数の要素を入れることができ、リストの要素には任意のタイプのオブジェクトを使用できます。リストは配列と考えることもできます。リスト内の要素の数は、要素を追加、削除、または置換する際に増加または減少します。

例

[]	空のリスト。
[1]	単一の要素 (整数) を含むリスト。
["Mike", 10, "Don", 20]	4 つの要素 (2 つの文字列要素と、2 つの整数要素) を含むリスト。
[[], [7], [8, 9]]	リストを含むリスト。各サブリストは、空のリスト、または整数要素のリスト。

```
x = 7; y = 2; z = 3;
[1, x, y, x + y]
```

整数のリスト。この例は、変数と式の使い方を示しています。

リストを変数に割り当てることができます。例えば、以下のようにします。

```
mylist1 = ["one", "two", "three"]
```

その後、このリストの特定の要素にアクセスできます。例えば、以下のようにします。

```
mylist[0]
```

これは以下のような出力になります。

```
one
```

大括弧 ([]) 内の数値は、インデックスと呼ばれ、リストの特定の要素を参照します。リストの各要素には、0 から始まるインデックスが付けられます。

1 つのリストの複数の要素の範囲を選択することもできます。これはスライスと呼ばれます。例えば、`x[1:3]` は、`x` の 2 番目と 3 番目のエレメントを選択します。終了インデックスは、選択を 1 つ過ぎています。

ストリング

文字列は、値として扱われる一連の不変の文字です。文字列は、新しい文字列になるすべての不変のシーケンス関数および演算子をサポートします。例えば、`"abcdef"[1:4]` は、`"bcd"` という出力になります。

Python では、文字は長さが 1 の文字列として表されます。

文字列リテラルは、単一引用符または三重引用符によって定義されます。単一引用符を使用して定義される文字列は行をまたぐことはできませんが、三重引用符を使用して定義される文字列は行をまたぐことができます。ストリングは、単一引用符 (') または二重引用符 (") で囲むことができます。引用文字には、エスケープされていない他の引用文字またはエスケープされた引用文字 (円記号 \) 文字が前に付きます) を含めることができます。

例

```
"This is a string"
'This is also a string'
'It's a string"
'This book is called "Python Scripting and Automation Guide".'
"This is an escape quote (\") in a quoted string"
```

空白文字で区切られた複数の文字列は、Python パーサーによって自動的に連結されます。これにより、長い文字列を入力したり、単一文字列で異なる種類の引用符を混在させたりすることができます。

```
"This string uses ' and " 'that string uses "."
```

これにより、次のように出力されます。

```
This string uses ' and that string uses ".
```

文字列は、いくつかの有用なメソッドをサポートしています。次の表に、これらのメソッドの一部を示します。

表 6. 文字列メソッド	
方法	使用法
<code>s.capitalize()</code>	<code>s</code> の頭文字を大文字にします。

表 6. 文字列メソッド (続き)	
方法	使用法
<code>s.count(ss {,start {,end}})</code>	<code>s[start:end]</code> 内の <code>ss</code> の出現回数をカウントします。
<code>s.startswith(str {, start {, end}})</code> <code>s.endswith(str {, start {, end}})</code>	<code>s</code> が <code>str</code> で始まっているかどうかをテストします。 <code>s</code> が <code>str</code> で終わっているかどうかをテストします。
<code>s.expandtabs({size})</code>	タブをスペース (デフォルトの <code>size</code> は 8) で置換します。
<code>s.find(str {, start {, end}})</code> <code>s.rfind(str {, start {, end}})</code>	<code>s</code> の中で <code>str</code> の最初のインデックスを検索します。見つからない場合、結果は -1 になります。 <code>rfind</code> は、右から左に検索します。
<code>s.index(str {, start {, end}})</code> <code>s.rindex(str {, start {, end}})</code>	<code>s</code> の中で <code>str</code> の最初のインデックスを検索します。見つからない場合、 <code>ValueError</code> が発生します。 <code>rindex</code> は、右から左に検索します。
<code>s.isalnum</code>	文字列が英数字かどうかを確認するためのテスト。
<code>s.isalpha</code>	文字列が英字かどうかを確認するためのテスト。
<code>s.isnum</code>	文字列が数値かどうかを確認するためのテスト。
<code>s.isupper</code>	文字列がすべて大文字かどうかを確認するためのテスト。
<code>s.islower</code>	文字列がすべて小文字かどうかを確認するためのテスト。
<code>s.isspace</code>	文字列がすべて空白文字かどうかを確認するためのテスト。
<code>s.istitle</code>	文字列が頭文字が大文字の一連の英数字かどうかを確認するためのテスト。
<code>s.lower()</code> <code>s.upper()</code> <code>s.swapcase()</code> <code>s.title()</code>	すべて小文字に変換します。 すべて大文字に変換します。 大文字/小文字をすべて逆に変換します。 すべてタイトル・ケースに変換します。
<code>s.join(seq)</code>	<code>seq</code> 内の文字列を <code>s</code> を区切り文字として結合します。
<code>s.splitlines({keep})</code>	<code>s</code> を複数行に分割します。 <code>keep</code> が <code>true</code> の場合、改行を保持します。
<code>s.split({sep {, max}})</code>	<code>s</code> を <code>sep</code> (デフォルトの <code>sep</code> は空白文字です) を使用して <code>max</code> 回まで「単語」に分割します。
<code>s.ljust(width)</code> <code>s.rjust(width)</code> <code>s.center(width)</code> <code>s.zfill(width)</code>	幅が <code>width</code> のフィールド内で文字列を左揃えします。 幅が <code>width</code> のフィールド内で文字列を右揃えします。 幅が <code>width</code> のフィールド内で文字列を中央揃えします。 0 で埋めます。

表 6. 文字列メソッド (続き)	
方法	使用法
<code>s.lstrip()</code> <code>s.rstrip()</code> <code>s.strip()</code>	先頭の空白文字を削除します。 末尾の空白文字を削除します。 先頭と末尾の空白文字を削除します。
<code>s.translate(str {,delc})</code>	<code>delc</code> 内の文字を削除した後、テーブルを使用して <code>s</code> を変換します。 <code>str</code> は、長さが <code>== 256</code> のストリングでなければなりません。
<code>s.replace(old, new {, max})</code>	文字列 <code>old</code> をすべて、または <code>max</code> 個の出現箇所を文字列 <code>new</code> で置き換えます。

注釈

注釈は、ポンド (またはハッシュ) 記号 (#) によって挿入されるコメントです。同じ行のポンド記号に続くすべてのテキストは、注釈の一部と見なされ、無視されます。注釈は、任意の桁から開始できます。以下の例で、注釈の使用法を示します。

```
#The HelloWorld application is one of the most simple
print 'Hello World' # print the Hello World line
```

ステートメントの構文

Python のステートメントのシンタックスは非常に単純です。一般に、各ソース行は単一ステートメントです。 `expression` および `assignment` ステートメントを除き、各ステートメントは `if` または `for` などのキーワード名によって導入されます。ブランク行または注釈行は、コード内の任意のステートメントの間の任意の場所に挿入できます。1 行に 2 つ以上のステートメントがある場合、各ステートメントをセミコロン (;) で区切る必要があります。

長いステートメントは、複数の行に続けることができます。この場合、次の行に継続するステートメントは、円記号 (\) で終了する必要があります。以下に例を示します。

```
x = "A loooooooooooooooooooooooooong string" + \
    "another loooooooooooooooooooooooooong string"
```

ある構造が括弧 (()), 大括弧 ([]), または中括弧 ({}) で囲まれている場合は、円記号を挿入することなく、ステートメントをカンマの後ろで新しい行に続けることができます。例えば、以下のようになります。

```
x = (1, 2, 3, "hello",
    "goodbye", 4, 5, 6)
```

識別子

識別子は、変数、関数、クラス、およびキーワードに名前を付けるために使用します。ID は任意の長さにすることができますが、先頭は大文字または小文字の英字、または下線文字 (_) にする必要があります。下線で始まる名前は、通常、内部名または私用名として予約されています。識別子の先頭文字の後ろに、英字、0 から 9 の数字、および下線文字をいくつでも自由に組み合わせて使用できます。

Jython には、変数、関数、またはクラスの名前に使用できない予約語がいくつかあります。これらの予約語は、以下のカテゴリーに分かれています。

- **ステートメント接頭部:** `assert`, `break`, `class`, `continue`, `def`, `del`, `elif`, `else`, `except`, `exec`, `finally`, `for`, `from`, `global`, `if`, `import`, `pass`, `print`, `raise`, `return`, `try`, および `while`
- **パラメーター接頭部:** `as`, `import`, および `in`

- 演算子: and、in、is、lambda、not、および or

不適切なキーワードを使用すると、通常 `SyntaxError` が発生します。

コードのブロック

コードのブロックは、単一ステートメントが期待される場所に使用されるステートメントのグループです。コードのブロックは、`if`、`elif`、`else`、`for`、`while`、`try`、`except`、`def`、および `class` のいずれのステートメントの後ろにも置くことができます。これらのステートメントの後ろにコロン (`:`) を使用して、コードのブロックを続けます。例えば、以下のようにします。

```
if x == 1:
    y = 2
    z = 3
elif:
    y = 4
    z = 5
```

コード・ブロックを区切るためにインデントが使用されます (Java では中括弧が使用される)。1つのブロック内のすべての行を同じ位置にインデントする必要があります。これは、インデントの変更が、コード・ブロックの終了を示すためです。通常は、レベルごとに4つのスペースでインデントします。行のインデントには、タブではなくスペースを使用することが推奨されています。スペースとタブを混在させることはできません。モジュールの最外部のブロックの行は、1桁目から開始する必要があります。そうでないと、`SyntaxError` が発生します。

1つのコード・ブロックを構成する複数のステートメント (コロンに続ける) は、セミコロンで区切って1行にすることもできます。例えば、以下のようにします。

```
if x == 1: y = 2; z = 3;
```

スクリプトへの引数の引き渡し

スクリプトに引数を渡すことは、変更せずにスクリプトを繰り返し使用できるため便利です。コマンド・ライン行で渡される引数は、リスト `sys.argv` 内の値として渡されます。渡される値の数は、コマンド `len(sys.argv)` を使用して取得できます。以下に例を示します。

```
import sys
print "test1"
print sys.argv[0]
print sys.argv[1]
print len(sys.argv)
```

この例では、`import` コマンドは、`sys` クラス全体をインポートして、このクラスに存在しているメソッド (`argv` など) を使用できるようにします。

この例のスクリプトは、以下の行を使用して起動できます。

```
/u/mjloos/test1 mike don
```

結果は以下の出力になります。

```
/u/mjloos/test1 mike don
test1
mike
don
3
```

例

print キーワードは、このキーワードの直後の引数を表示します。ステートメントの後ろにコンマを続けると、改行は出力に含まれません。以下に例を示します。

```
print "This demonstrates the use of a",
print " comma at the end of a print statement."
```

これは以下のような出力になります。

```
This demonstrates the use of a comma at the end of a print statement.
```

for ステートメントは、コードのブロックを反復するために使用します。以下に例を示します。

```
mylist1 = ["one", "two", "three"]
for lv in mylist1:
    print lv
    continue
```

この例では、3つの文字列がリスト mylist1 に割り当てられます。リストの各要素が1行に1つずつ出力されます。これは以下のような出力になります。

```
one
two
three
```

この例では、for ループが要素ごとのコード・ブロックを実装するたびに、イテレーター lv がリスト mylist1 の各要素の値を順にとります。イテレーターは、任意の長さの有効な ID にすることができます。

if ステートメントは、条件ステートメントです。条件を評価し、評価の結果に基づいて true または false を返します。以下に例を示します。

```
mylist1 = ["one", "two", "three"]
for lv in mylist1:
    if lv == "two":
        print "The value of lv is ", lv
    else:
        print "The value of lv is not two, but ", lv
    continue
```

この例では、イテレーター lv の値が評価されます。lv の値が two の場合は、lv が two でない場合に返される文字列とは異なる文字列が返されます。これにより、次のように出力されます。

```
The value of lv is not two, but one
The value of lv is two
The value of lv is not two, but three
```

数学メソッド

math モジュールから、有用な数学メソッドにアクセスできます。次の表に、これらのメソッドの一部を示します。特に指定のない限り、すべての値は浮動小数点として返されます。

表 7. 数学メソッド	
方法	使用法
math.ceil(x)	x の天井値を浮動小数点として返します。これは、x 以上の最小の整数です。
math.copysign(x, y)	y の符号を付けて x を返します。copysign(1, -0.0) 戻り値 -1
math.fabs(x)	x の絶対値を返します。

表 7. 数学メソッド (続き)	
方法	使用法
<code>math.factorial(x)</code>	x 階乗を返します。x が負の場合、または整数でない場合、 <code>ValueError</code> が発生します。
<code>math.floor(x)</code>	x の床値を浮動小数点として返します。これは、x 以下の最大の整数です。
<code>math.frexp(x)</code>	x の仮数 (m) と指数 (e) をペア (m, e) として返します。m は浮動小数点、e は整数で、 $x == m * 2^{**e}$ となります。x がゼロの場合は (0.0, 0) を返し、それ以外の場合は $0.5 \leq \text{abs}(m) < 1$ を返します。
<code>math.fsum(iterable)</code>	iterable の中の値の正確な浮動小数点の和を返します。
<code>math.isinf(x)</code>	浮動小数点 x が正または負の無限大かどうかをチェックします。
<code>math.isnan(x)</code>	浮動小数点 x が NaN (非数値) かどうかをチェックします。
<code>math.ldexp(x, i)</code>	$x * (2^{**i})$ を返します。これは、本質的に関数 <code>frexp</code> の逆です。
<code>math.modf(x)</code>	x の小数部と整数部を返します。結果は両方とも x の符号を引き継ぎ、浮動小数点です。
<code>math.trunc(x)</code>	Integral に切り捨てられた Real 値 x を返します。
<code>math.exp(x)</code>	e^{**x} を返します。
<code>math.log(x[, base])</code>	指定した値 base に対する x の対数を返します。base を指定しない場合は、x の自然対数が返されます。
<code>math.log1p(x)</code>	$1+x$ (base e) の自然対数を返します。
<code>math.log10(x)</code>	x の 10 を底とする対数を返します。
<code>math.pow(x, y)</code>	x の累乗 y を返します。pow(1.0, x) および pow(x, 0.0) は、x がゼロまたは NaN であっても、常に 1 を返します。
<code>math.sqrt(x)</code>	x の平方根を返します。

数学関数に加えて、有用な三角関数メソッドもあります。次の表に、これらのメソッドを示します。

表 8. 三角関数メソッド	
方法	使用法
<code>math.acos(x)</code>	x の逆余弦をラジアンで返します。
<code>math.asin(x)</code>	x の逆正弦をラジアンで返します。
<code>math.atan(x)</code>	x の逆正接をラジアンで返します。
<code>math.atan2(y, x)</code>	$\text{atan}(y / x)$ をラジアンで返します。

表 8. 三角関数メソッド (続き)	
方法	使用法
<code>math.cos(x)</code>	x の余弦をラジアンで返します。
<code>math.hypot(x, y)</code>	ユークリッド・ノルム <code>sqrt(x*x + y*y)</code> を返します。これは原点から点 (x, y) へのベクトルの長さです。
<code>math.sin(x)</code>	x の正弦をラジアンで返します。
<code>math.tan(x)</code>	x の正接をラジアンで返します。
<code>math.degrees(x)</code>	角 x をラジアンから度に変換します。
<code>math.radians(x)</code>	角 x を度からラジアンに変換します。
<code>math.acosh(x)</code>	x の逆双曲線余弦を返します。
<code>math.asinh(x)</code>	x の逆双曲線正弦を返します。
<code>math.atanh(x)</code>	x の逆双曲線正接を返します。
<code>math.cosh(x)</code>	x の双曲線余弦を返します。
<code>math.sinh(x)</code>	x の双曲線余弦を返します。
<code>math.tanh(x)</code>	x の双曲線正接を返します。

2 つの数学定数もあります。 `math.pi` の値は、数学定数 pi です。 `math.e` の値は、数学定数 e です。

非 ASCII 文字の使用

非 ASCII 文字を使用するには、Python では、文字列を Unicode に明示的にエンコードまたはデコードする必要があります。IBM スポス モデラー では、Python スクリプトは UTF-8 (非 ASCII 文字をサポートする標準 Unicode) でエンコードされていると想定されます。以下のスクリプトは、Python コンパイラーがスポス モデラー によって UTF-8 に設定されているため、コンパイルされます。

```
stream = modeler.script.stream()
filenode = stream.createAt("variablefile", "テストノード", 96, 64)
```

しかし、結果ノードのラベルは正しくありません。



図 3. 非 ASCII 文字を含むノード・ラベル (正しく表示されていない)

ストリング・リテラル自体が Python によって ASCII 文字列に変換されているため、このラベルは正しくありません。

Python では、文字列リテラルの前に u 文字を追加することによって、Unicode 文字列リテラルを指定できます。

```
stream = modeler.script.stream()
filenode = stream.createAt("variablefile", u"テストノード", 96, 64)
```

これにより、Unicode 文字列が作成され、ラベルが正しく表示されます。



テストノード

図 4. 非 ASCII 文字を含むノード・ラベル (正しく表示されている)

Python と Unicode の使用は、本書の範囲を超えた大きなトピックです。このトピックを詳細に扱った書籍やオンライン情報源が数多くあります。

オブジェクト指向プログラミング

オブジェクト指向プログラミングは、対象問題のモデルをプログラム内で作成するという概念に基づいています。オブジェクト指向プログラミングにより、プログラミング・エラーが減り、コードの再使用が促進されます。Python は、オブジェクト指向言語です。Python で定義されるオブジェクトには、以下の特徴があります。

- **恒等式:** 各オブジェクトは個別であり、これはテスト可能でなければなりません。is テストと is not テストは、この目的のために存在しています。
- **状態:** 各オブジェクトは、状態を格納できる必要があります。フィールドやインスタンス変数などの属性は、この目的のために存在しています。
- **振る舞い:** 各オブジェクトは、状態を操作できる必要があります。メソッドは、この目的のために存在します。

Python には、オブジェクト指向プログラミングをサポートするための以下の特徴があります。

- **クラス・ベースのオブジェクト作成:** クラスは、オブジェクトを作成するためのテンプレートです。オブジェクトは、振る舞いが関連づけられているデータ構造です。
- **ポリモアフィズムによる継承:** Python は、単一継承と多重継承をサポートしています。Python のすべてのインスタンス・メソッドは、ポリモアフィックであり、サブクラスによるオーバーライドが可能です。
- **データ隠蔽によるカプセル化:** Python では、属性を隠すことができます。隠すと、クラスの外側からは、そのクラスのメソッドによってのみ属性にアクセスできるようになります。クラスには、データを変更するためのメソッドを実装します。

クラスの定義

Python クラスの中では、変数とメソッドの両方を定義できます。Java と異なり、Python では、1つのソース・ファイル(モジュール)で任意の数の公開クラスを定義できます。したがって、Python のモジュールは Java のパッケージに似ていると考えることができます。

Python では class ステートメントを使用してクラスを定義します。class ステートメントは、次の形式になっています。

```
class name (superclasses): statement
```

または

```
class name (superclasses):  
    assignment  
    :  
    :  
    function  
    :  
    :
```

クラスを定義するときには、任意の数の代入 ステートメントを記述することができます (記述しなくても構いません)。これにより、クラスのすべてのインスタンスで共有されるクラス属性が作成されます。また、任意の数の関数 定義を記述することもできます (記述しなくても構いません)。これらの関数定義により、メソッドが作成されます。スーパークラスのリストはオプションです。

クラス名はスコープの中 (モジュール、関数、またはクラスの中) で固有でなければなりません。複数の変数を定義して同じクラスを参照することができます。

クラス・インスタンスの作成

クラスは、クラス (共有) 属性の保持やクラス・インスタンスの作成に使用します。クラスのインスタンスを作成するには、そのクラスが関数であるかのように呼び出します。たとえば、次のクラスを考慮してください。

```
class MyClass:
    pass
```

クラスを完結させるためにはステートメントが必要ですがプログラムとしては動作が不要であるため、ここでは `pass` ステートメントを使用しています。

以下のステートメントは、クラス `MyClass` のインスタンスを作成します。

```
x = MyClass()
```

クラス・インスタンスへの属性の追加

Java と異なり、Python ではクライアントがクラスのインスタンスに属性を追加することができます。変更されるインスタンスは 1 つだけです。例えば、インスタンス `x` に複数の属性を追加するには、以下のようにしてそのインスタンスに新しい値を設定します。

```
x.attr1 = 1
x.attr2 = 2
.
.
x.attrN = n
```

クラス属性およびメソッドの定義

クラスにバインドされた変数はすべてクラス属性です。クラス内で定義された関数はすべてメソッドです。メソッドは、クラスのインスタンス (慣習として `self` と呼びます) を第 1 引数として受け取ります。例えば、クラス属性およびメソッドを定義するには、以下のコードを入力します。

```
class MyClass
    attr1 = 10          #class attributes
    attr2 = "hello"

    def method1(self):
        print MyClass.attr1    #reference the class attribute

    def method2(self):
        print MyClass.attr2    #reference the class attribute

    def method3(self, text):
        self.text = text       #instance attribute
        print text, self.text  #print my argument and my attribute

    method4 = method3          #make an alias for method3
```

クラスの内側では、クラス属性に対するすべての参照をクラス名で修飾する必要があります (`MyClass.attr1` など)。インスタンス属性に対する参照は、すべて `self` 変数で修飾する必要があります (`self.text` など)。クラスの外側では、クラス属性に対するすべての参照をクラス名で修飾するか (`MyClass.attr1` など)、クラスのインスタンスで修飾する (`x` をクラスのインスタンスとすると `x.attr1` などとする) 必要があります。クラスの外側では、インスタンス変数に対するすべての参照をクラスのインスタンスで修飾する必要があります (`x.text` など)。

非表示変数

プライベート変数を作成することにより、データを隠蔽することができます。プライベート変数にアクセスできるのはそのクラス自体に限られます。__xxx または __xxx_yyy という形式で (2 個の下線を前に付けて) 名前を宣言すると、Python パーサーは、宣言された名前に自動的にクラス名を追加して隠蔽された変数を作成します。例を示します。

```
class MyClass:
    __attr = 10    #private class attribute

    def method1(self):
        pass

    def method2(self, p1, p2):
        pass

    def __privateMethod(self, text):
        self.__text = text    #private attribute
```

Java と異なり、Python では、インスタンス変数に対する参照はすべて **self** で修飾する必要があります。暗黙的な **this** の使用はありません。

継承

クラスを継承する機能は、オブジェクト指向プログラミングの根幹をなします。Python は、単一継承と多重継承の両方をサポートしています。単一継承は、スーパークラスが 1 つしか存在できないことを意味します。多重継承は、複数のスーパークラスが存在できることを意味します。

継承は、他のクラスのサブクラスを定義することで実装します。任意の数の Python クラスをスーパークラスにすることができます。Python の Jython 実装では、直接または間接に継承できる Java クラスは 1 つだけです。スーパークラスを提供する必要はありません。

スーパークラスのすべての属性やメソッドはいずれのサブクラスにも存在し、そのクラス自体によって使用できるほか、属性やメソッドが隠蔽されていなければ任意のクライアントから使用することもできます。サブクラスのインスタンスは任意の場所で使用でき、スーパークラスのインスタンスも使用できます。これがポリモフィズムの一例です。これらの機能によって再利用が可能になり、拡張が容易になります。

例

```
class Class1: pass    #no inheritance

class Class2: pass

class Class3(Class1): pass    #single inheritance

class Class4(Class3, Class2): pass    #multiple inheritance
```

第 3 章 IBM SPSS Modeler でのスクリプト

スクリプトの種類

IBM SPSS Modeler には、以下の 3 種類のスクリプトがあります。

- ストリーム・スクリプト は、単一ストリームの実行を制御するために使用され、ストリーム内に格納されます。
- スーパーノード・スクリプト は、スーパーノードの動作を制御するために使用されます。
- スタンドアロン スクリプトまたはセッション・スクリプト は、さまざまなストリームにわたって実行を調整するために使用できます。

さまざまなメソッドを IBM SPSS Modeler のスクリプトで使用することができ、これらメソッドによって SPSS Modeler の広範な機能にアクセスできます。これらのメソッドは、より高度な機能を作成するために 39 ページの『第 4 章 スクリプト API』でも使用されます。

ストリーム、スーパーノード・ストリーム、およびダイアグラム

多くの場合、ストリーム という語は、ファイルからロードされるストリームであれ、スーパーノード内で使用されるストリームであれ、同じ意味を持ちます。一般に、ストリームは、互いに接続された実行可能なノードの集合を意味します。しかし、スクリプトの場合は、あらゆる場所ですべての操作がサポートされるわけではありません。つまり、スクリプト作成者は、どのストリーム・バリエントを使用しているのかを認識している必要があります。

ストリーム

ストリームは、IBM SPSS Modeler の主なドキュメント・タイプです。ストリームは保存、ロード、編集、および実行することができます。ストリームには、パラメーター、グローバル値、スクリプト、およびその他の情報を関連付けることもできます。

スーパーノード・ストリーム

スーパーノード・ストリーム は、スーパーノード内で使用される種類のストリームです。通常のストリームと同様、互いにリンクされているノードが含まれています。スーパーノード・ストリームは、以下の点で通常のストリームと異なっています。

- パラメーターおよびスクリプトは、スーパーノード・ストリームではなく、スーパーノード・ストリームを所有しているスーパーノードに関連付けられています。
- スーパーノード・ストリームには、スーパーノードの種類に応じて、追加の入力コネクター・ノードや出力コネクター・ノードがあります。これらのコネクター・ノードは、スーパーノード・ストリームに情報を渡したり、スーパーノード・ストリームから情報を取り出したりするために使用され、スーパーノードの作成時に自動的に作成されます。

ダイアグラム

ダイアグラム という用語は、通常のストリームとスーパーノード・ストリームの両方でサポートされる機能 (ノードの追加や削除、ノード間の接続の変更など) を含んでいます。

ストリームの実行

以下の例は、ストリーム内のすべての実行可能ノードを実行する最もシンプルなタイプのストリーム・スクリプトです。

```
modeler.script.stream().runAll(None)
```

以下の例も、ストリーム内のすべての実行可能ノードを実行します。

```
stream = modeler.script.stream()
stream.runAll(None)
```

この例では、ストリームを変数 `stream` に格納しています。通常、スクリプトはストリームまたはストリーム内のノードを変更するために使用されるため、ストリームを変数に格納すると便利です。ストリームを格納する変数を作成することによって、スクリプトはより簡潔になります。

スクリプト・コンテキスト

`modeler.script` モジュールは、スクリプトが実行されるコンテキストを提供します。このモジュールは、実行時に SPSS Modeler スクリプトに自動的にインポートされます。このモジュールは、スクリプトがその実行環境にアクセスするための方法を提供する 4 つの関数を定義しています。

- `session()` 関数は、スクリプトのセッションを返します。セッションは、ストリームを実行するために使用されているロケールや、SPSS Modeler バックエンド (ローカル・プロセス、またはネットワーク SPSS Modeler Server) などの情報を定義します。
- `stream()` 関数は、ストリームとスーパーノード・スクリプトで使用できます。この関数は、実行中のストリーム・スクリプトまたはスーパーノード・スクリプトを所有しているストリームを返します。
- `diagram()` 関数は、スーパーノード・スクリプトで使用できます。この関数は、スーパーノード内のダイアグラムを返します。その他のスクリプトのタイプの場合、この関数は `stream()` 関数と同じ内容を返します。
- `supernode()` 関数は、スーパーノード・スクリプトで使用できます。この関数は、実行中のスクリプトを所有しているスーパーノードを返します。

これら 4 つの関数と出力を次の表に要約します。

表 9. <code>modeler.script</code> 関数の要約				
スクリプト・タイプ	<code>session()</code>	<code>stream()</code>	<code>diagram()</code>	<code>supernode()</code>
スタンドアロン	セッションを返します	スクリプト起動時の現在の管理対象ストリーム (例えば、バッチ・モード <code>-stream</code> オプションによって渡されたストリーム) か、 <code>None</code> を返します。	<code>stream()</code> と同じ	適用外
ストリーム	セッションを返します	ストリームを返します	<code>stream()</code> と同じ	適用外
スーパーノード	セッションを返します	ストリームを返します	スーパーノード・ストリームを返します	スーパーノードを返します

`modeler.script` モジュールは、終了コードでスクリプトを終了する方法も定義します。`exit(exit-code)` 関数は、スクリプトの実行を停止し、指定された整数の終了コードを返します。

ストリーム用に定義されているメソッドの 1 つに `runAll(List)` があります。このメソッドは、すべての実行可能ノードを実行します。ノードを実行することで生成されるモデルまたは出力は、指定されたリストに追加されます。

通常、ストリームを実行すると、モデルやグラフなどの出力が生成されます。この出力をキャプチャーするために、スクリプトは、リストに初期化される変数を提供できます。例えば、以下のとおりです。

```
stream = modeler.script.stream()
results = []
stream.runAll(results)
```

実行が完了すると、実行によって生成されたオブジェクトに **results** リストからアクセスできます。

既存のノードの参照

多くの場合、ストリームは、ストリームの実行前に変更する必要があるいくつかのパラメーターを使用して事前構築されています。これらのパラメーターを変更するには、以下の作業を行います。

1. 関連するストリーム内のノードを見つける。
2. ノードまたはストリーム (あるいは両方) の設定を変更する。

ノードの検索

ストリームでは、さまざまな方法で既存のノードを見つけることができます。これらのメソッドを次の表に要約します。

表 10. 既存のノードを見つけるためのメソッド		
方法	戻り値の型	説明
<code>s.findAll(type, label)</code>	コレクション	指定したデータ型とラベルを持つすべてのノードのリストを返します。データ型またはラベルのいずれかが None の場合は、もう一方のパラメーターが使用されます。
<code>s.findAll(filter, recursive)</code>	コレクション	指定したフィルターで受け入れられるすべてのノードの集合を返します。 <code>recursive</code> フラグが True の場合は、指定したストリーム内のスーパーノードも検索されます。
<code>s.findById(id)</code>	ノード	指定した ID のノードを返すか、そのようなノードが存在しない場合は None を返します。検索は現行ストリームに限定されます。
<code>s.findByType(type, label)</code>	ノード	指定したデータ型またはラベルを持つノード、あるいはその両方を持つノードを返します。データ型または名前のいずれかが None の場合は、もう一方のパラメーターが使用されます。一致するノードが複数ある場合は、任意のノードが返されます。一致するノードがない場合、戻り値は None です。

表 10. 既存のノードを見つけるためのメソッド (続き)

方法	戻り値の型	説明
<code>s.findDownstream(fromNodes)</code>	コレクション	指定したノードのリストから検索し、指定したノードの下流にある一連のノードを返します。返されるリストには、最初に指定したノードも含まれます。
<code>s.findUpstream(fromNodes)</code>	コレクション	指定したノードのリストから検索し、指定したノードの上流にある一連のノードを返します。返されるリストには、最初に指定したノードも含まれます。
<code>s.findProcessorForID(String id, boolean recursive)</code>	ノード	指定した ID のノードを返すか、そのようなノードが存在しない場合は <code>None</code> を返します。recursive フラグが <code>true</code> の場合は、このダイアグラムの中の複合ノードも検索されます。

例えば、スクリプトがアクセスする必要がある単一のフィルター・ノードがストリームに含まれている場合、そのフィルター・ノードは、以下のスクリプトを使用して見つけることができます。

```
stream = modeler.script.stream()
node = stream.findByType("filter", None)
...
```

あるいは、ノードの ID (ノード・ダイアログ・ボックスの「注釈」タブに示されている) が分かる場合は、その ID を使用してノードを検索できます。例えば、以下のようになります。

```
stream = modeler.script.stream()
node = stream.findByID("id32FJT71G2") # the filter node ID
...
```

プロパティを設定する

ノード、ストリーム、モデル、および出力のすべてには、アクセス可能で、ほとんどの場合に設定可能なプロパティがあります。通常、プロパティは、オブジェクトの動作および外観を変更するために使用されます。オブジェクトのプロパティのアクセスおよび設定に使用できるメソッドを次の表に要約します。

表 11. オブジェクトのプロパティのアクセスおよび設定のためのメソッド

方法	戻り値の型	説明
<code>p.getPropertyValue(propertyName)</code>	オブジェクト	指定したプロパティの値を返すか、そのようなプロパティが存在しない場合は <code>None</code> を返します。
<code>p.setPropertyValue(propertyName, value)</code>	適用外	指定したプロパティの値を設定します。

表 11. オブジェクトのプロパティのアクセスおよび設定のためのメソッド (続き)		
方法	戻り値の型	説明
<code>p.setPropertyValues(properties)</code>	適用外	指定したプロパティの値を設定します。プロパティ・マップの各項目は、プロパティ名を表すキーと、そのプロパティに割り当てる必要がある値で構成されています。
<code>p.getKeyedPropertyValue(propertyName, keyName)</code>	オブジェクト	指定したプロパティの値および関連付けられているキーを返すか、そのようなプロパティまたはキーが存在しない場合は <code>None</code> を返します。
<code>p.setKeyedPropertyValue(propertyName, keyName, value)</code>	適用外	指定したプロパティおよびキーの値を設定します。

例えば、ストリームの先頭にある可変長ファイル・ノードの値を設定する場合は、以下のスクリプトを使用できます。

```
stream = modeler.script.stream()
node = stream.findByType("variablefile", None)
node.setPropertyValue("full_filename", "$CLEO/DEMOS/DRUG1n")
...
```

あるいは、フィルター・ノードからフィールドをフィルタリングできます。この場合は、フィールド名に対して値も入力します。例えば、以下のようになります。

```
stream = modeler.script.stream()
# Locate the filter node ...
node = stream.findByType("filter", None)
# ... and filter out the "Na" field
node.setKeyedPropertyValue("include", "Na", False)
```

ノードの作成とストリームの変更

新しいノードを既存のストリームに追加する場合があります。既存のストリームにノードを追加するには、通常以下の作業を行います。

1. ノードを作成する。
2. ノードを既存のストリーム・フローにリンクする。

ノードの作成

ストリームでは、さまざまな方法でノードを作成できます。これらのメソッドを次の表に要約します。

表 12. ノードを作成するためのメソッド		
方法	戻り値の型	説明
<code>s.create(nodeType, name)</code>	ノード	指定したデータ型のノードを作成して、指定したストリームに追加します。

表 12. ノードを作成するためのメソッド (続き)		
方法	戻り値の型	説明
<code>s.createAt(nodeType, name, x, y)</code>	ノード	指定したデータ型のノードを作成して、指定したストリームの指定した場所に追加します。 $x < 0$ または $y < 0$ の場合、場所は設定されません。
<code>s.createModelApplier(modelOutput, name)</code>	ノード	提供されたモデル出力オブジェクトから派生したモデル・アプライヤー・ノードを作成します。

例えば、ストリーム内に新しいデータ型ノードを作成するには、以下のスクリプトを使用できます。

```
stream = modeler.script.stream()
# Create a new type node
node = stream.create("type", "My Type")
```

ノードのリンクとリンク解除

ストリーム内に新しいノードを作成する場合、そのノードを使用するにはノードのシーケンスに接続する必要があります。ストリームには、ノードをリンクおよびリンク解除するための多くのメソッドがあります。これらのメソッドを次の表に要約します。

表 13. ノードをリンクおよびリンク解除するためのメソッド		
方法	戻り値の型	説明
<code>s.link(source, target)</code>	適用外	入力ノードとターゲット・ノードの間に新しいリンクを作成します。
<code>s.link(source, targets)</code>	適用外	入力ノードと指定されたリスト内の各ターゲットの間に新しいリンクを作成します。
<code>s.linkBetween(inserted, source, target)</code>	適用外	他の 2 つのノード・インスタンス (入力ノードとターゲット・ノード) の間にノードを接続し、挿入したノードの位置がこれらのノードの間になるように設定します。入力ノードとターゲット・ノードの間の直接リンクが最初に削除されます。
<code>s.linkPath(path)</code>	適用外	ノード・インスタンスの間の新しいパスを作成します。最初のノードが 2 番目のノードにリンクされ、2 番目のノードが 3 番目のノードにリンクされ、以下同様にリンクされます。
<code>s.unlink(source, target)</code>	適用外	入力ノードとターゲット・ノードの間の直接リンクを削除します。
<code>s.unlink(source, targets)</code>	適用外	入力ノードと指定されたターゲット・リスト内の各オブジェクトの間の直接リンクを削除します。

表 13. ノードをリンクおよびリンク解除するためのメソッド (続き)		
方法	戻り値の型	説明
<code>s.unlinkPath(path)</code>	適用外	ノード・インスタンスの間に存在するパスをすべて削除します。
<code>s.disconnect(node)</code>	適用外	指定されたノードと、指定したストリーム内の他のすべてのノードの間のリンクを削除します。
<code>s.isValidLink(source, target)</code>	<i>Boolean</i>	指定した入力ノードとターゲット・ノードの間にリンクを作成できる場合は <i>True</i> を返します。 このメソッドは、指定したストリームに両方のオブジェクトが属していること、入力ノードがリンクを提供でき、ターゲット・ノードがリンクを受け取れること、このようなリンクを作成してもストリーム内に循環が発生しないことをチェックします。

以下に示すサンプル・スクリプトは、以下の 5 つのタスクを実行します。

1. 可変長ファイル入力ノード、フィルター・ノード、およびテーブル出力ノードを作成する。
2. ノード同士を接続する。
3. 可変長ファイル入力ノードにファイル名を設定する。
4. 結果出力から「Drug」フィールドをフィルタリングする。
5. テーブル・ノードを実行する。

```
stream = modeler.script.stream()
filenode = stream.createAt("variablefile", "My File Input ", 96, 64)
filternode = stream.createAt("filter", "Filter", 192, 64)
tablenode = stream.createAt("table", "Table", 288, 64)
stream.link(filenode, filternode)
stream.link(filternode, tablenode)
filenode.setPropertyValue("full_filename", "$CLEO_DEMOS/DRUG1n")
filternode.setKeyedPropertyValue("include", "Drug", False)
results = []
tablenode.run(results)
```

ノードのインポート、置換、および削除

ノードの作成や接続だけでなく、多くの場合にストリームのノードの置換や削除も必要です。ノードのインポート、置換、および削除に使用できるメソッドを次の表に要約します。

表 14. ノードをインポート、置換、および削除するためのメソッド		
方法	戻り値の型	説明
<code>s.replace(originalNode, replacementNode, discardOriginal)</code>	適用外	指定したストリームの指定したノードを置換します。元のノードと置換ノードの両方が、指定したストリームによって所有されている必要があります。

表 14. ノードをインポート、置換、および削除するためのメソッド (続き)

方法	戻り値の型	説明
<code>s.insert(source, nodes, newIDs)</code>	リスト	指定されたリスト内のノードのコピーを挿入します。指定されたリスト内のすべてのノードが、指定したストリームに含まれていると想定されます。 <code>newIDs</code> フラグは、ノードごとに新しい ID を生成するのか、または既存の ID をコピーして使用するのかを示します。ストリーム内のすべてのノードの ID は固有であると想定されているため、指定したストリームとソース・ストリームが同じである場合、このフラグを <code>True</code> に設定する必要があります。このメソッドは新しく挿入されたノードのリストを返しますが、ノードの順序は定義されていません (つまり、順序は入力リストのノードの順序と必ずしも同じであるとは限りません)。
<code>s.delete(node)</code>	適用外	指定したストリームから指定したノードを削除します。ノードは、指定したストリームによって所有されている必要があります。
<code>s.deleteAll(nodes)</code>	適用外	指定したストリームから指定したすべてのノードを削除します。集合内のすべてのノードが、指定したストリームに属している必要があります。
<code>s.clear()</code>	適用外	指定したストリームからすべてのノードを削除します。

ストリーム内のノードのトラバース

一般的な要件として、特定のノードの上流または下流にあるノードを識別したい場合があります。ストリームには、これらのノードを識別するために使用できる多くのメソッドがあります。これらのメソッドを次の表に要約します。

表 15. 上流または下流のノードを識別するためのメソッド

方法	戻り値の型	説明
<code>s.iterator()</code>	イテレーター	指定したストリームに含まれているノード・オブジェクトのイテレーターを返します。 <code>next()</code> 関数の呼び出しの間にストリームが変更される場合、イテレーターの動作は未定義です。
<code>s.predecessorAt(node, index)</code>	ノード	指定したノードの指定された直接の先行ノードを返すか、インデックスが境界を超えている場合は <code>None</code> を返します。

表 15. 上流または下流のノードを識別するためのメソッド (続き)

方法	戻り値の型	説明
<code>s.predecessorCount(node)</code>	<i>int</i>	指定されたノードの直接の先行ノードの数を返します。
<code>s.predecessors(node)</code>	リスト	指定されたノードの直接の先行ノードを返します。
<code>s.successorAt(node, index)</code>	ノード	指定したノードの指定した直接の後続ノードを返すか、インデックスが境界を超えている場合は <code>None</code> を返します。
<code>s.successorCount(node)</code>	<i>int</i>	指定されたノードの直接の後続ノードの数を返します。
<code>s.successors(node)</code>	リスト	指定されたノードの直接の後続ノードを返します。

項目の消去または削除

従来のスクリプトでは、以下の例のような、`clear` コマンドのさまざまな使用法がサポートされています。

- `clear outputs` は、すべての出力項目をマネージャ パレットから削除します。
- `clear generated palette` は、「モデル」パレットからすべてのモデル ナゲットを消去します。
- `clear stream` は、ストリームの中身を削除します。

Python スクリプトでは、同様の関数セットがサポートされます。ストリーム マネージャ、出力マネージャ、およびモデル マネージャを消去するには、`removeAll()` コマンドを使用します。以下に例を示します。

- ストリーム マネージャを消去する場合:

```
session = modeler.script.session()
session.getStreamManager.removeAll()
```

- 出力マネージャを消去する場合:

```
session = modeler.script.session()
session.getDocumentOutputManager().removeAll()
```

- モデル マネージャを消去する場合:

```
session = modeler.script.session()
session.getModelOutputManager().removeAll()
```

ノードに関する情報の入手

ノードは、データ・インポート・ノードおよびデータ・エクスポート・ノード、モデル構築ノード、その他の種類のノードなど、さまざまなカテゴリーに分類されます。各ノードには、ノードに関する情報を見つげるために使用できる多くのメソッドがあります。

ノードの ID、名前、およびラベルを取得するために使用できるメソッドを次の表に要約します。

表 16. ノードの ID、名前、およびラベルを取得するためのメソッド

方法	戻り値の型	説明
<code>n.getLabel()</code>	<i>string</i>	指定したノードの表示ラベルを返します。ラベルがプロパティ <code>custom_name</code> の値となるのは、このプロパティが空文字列ではなく、 <code>use_custom_name</code> プロパティが設定されていない場合のみです。これ以外の場合、ラベルは <code>getName()</code> の値になります。
<code>n.setLabel(label)</code>	適用外	指定したノードの表示ラベルを設定します。新しいラベルが空文字列ではない場合、この文字列がプロパティ <code>custom_name</code> に割り当てられ、指定したラベルが優先されるようにプロパティ <code>use_custom_name</code> に <code>False</code> が割り当てられます。これ以外の場合は、空文字列が <code>custom_name</code> に割り当てられ、プロパティ <code>use_custom_name</code> に <code>True</code> が割り当てられます。
<code>n.getName()</code>	<i>string</i>	指定されたノードの名前を返します。
<code>n.getID()</code>	<i>string</i>	指定したノードの ID を返します。新しいノードが作成されるたびに、新しい ID が作成されます。この ID は、ストリームの一部としてノードが保存されるときに、ノードで永続化され、ストリームを開いたときにノード ID が保持されるようになります。ただし、保存したノードがストリームに挿入される場合、挿入されたノードは新しいオブジェクトと見なされ、新しい ID が割り当てられます。

ノードに関するその他の情報を取得するために使用できるメソッドを次の表に要約します。

表 17. ノードに関する情報を取得するためのメソッド

方法	戻り値の型	説明
<code>n.getTypeName()</code>	<i>string</i>	このノードのスクリプト名を返します。これは、このノードの新しいインスタンスを作成するために使用できる名前と同じです。
<code>n.isInitial()</code>	<i>Boolean</i>	これが最初のノード (ストリームの先頭にあるノード) である場合は、 <code>True</code> を返します。

表 17. ノードに関する情報を取得するためのメソッド (続き)

方法	戻り値の型	説明
<code>n.isInline()</code>	<i>Boolean</i>	これがインライン・ノード (ストリームの中間にあるノード) である場合は、 True を返します。
<code>n.isTerminal()</code>	<i>Boolean</i>	これが終端 ノード (ストリームの末尾にあるノード) である場合は、 True を返します。
<code>n.getXPosition()</code>	<i>int</i>	ストリーム内のノードの x 位置オフセットを返します。
<code>n.getYPosition()</code>	<i>int</i>	ストリーム内のノードの y 位置オフセットを返します。
<code>n.setXYPosition(x, y)</code>	適用外	ストリーム内のノードの位置を設定します。
<code>n.setPositionBetween(source, target)</code>	適用外	指定されたノードの間に位置するようにストリーム内のノードの位置を設定します。
<code>n.isCacheEnabled()</code>	<i>Boolean</i>	キャッシュが有効な場合は True を返し、そうでない場合は False を返します。
<code>n.setCacheEnabled(val)</code>	適用外	このオブジェクトのキャッシュを有効または無効にします。キャッシュがいっぱいの場合にキャッシュが無効になると、キャッシュはフラッシュされます。
<code>n.isCacheFull()</code>	<i>Boolean</i>	キャッシュがいっぱいの場合は True を返し、そうでない場合は False を返します。
<code>n.flushCache()</code>	適用外	このノードのキャッシュをフラッシュします。キャッシュが有効でない場合やいっぱいでない場合、影響はありません。

第 4 章 スクリプト API

スクリプト API の概要

スクリプト API により、幅広い SPSS Modeler 機能にアクセスすることができます。ここまで説明してきたメソッドはいずれも API の一部であり、追加でインポートを行わなくてもスクリプト内から暗黙的にアクセスすることができます。ただし、API クラスを参照する必要がある場合は、以下のステートメントで明示的に API をインポートする必要があります。

```
import modeler.api
```

この import ステートメントは、多くのスクリプト API の例で必要になります。

スクリプト API を通じて使用可能なクラス、メソッド、およびパラメータの完全なガイドは、「*IBM SPSS Modeler Python Scripting API Reference Guide*」という文書に含まれています。

例 1: カスタム・フィルターを使用したノードの検索

29 ページの『ノードの検索』のセクションでは、検索基準としてノードのタイプ名を使用してストリームのノードを検索する例を示しました。場合によっては、より汎用的な検索が必要になります。そのような検索を実装するには、`NodeFilter` クラスおよびストリームの `findAll()` メソッドを使用します。この種の検索は以下の 2 段階で行います。

1. `NodeFilter` を拡張し、カスタム・バージョンの `accept()` メソッドを実装する新しいクラスを作成します。
2. この新しいクラスのインスタンスでストリームの `findAll()` メソッドを呼び出します。これにより、`accept()` メソッドで定義された基準を満たすすべてのノードが返されます。

ストリームのノードのうち、ノードのキャッシュが有効になっているノードを検索する方法を以下の例に示します。返されたノードのリストを使用して、それらのノードのキャッシュをフラッシュするか無効化することができます。

```
import modeler.api

class CacheFilter(modeler.api.NodeFilter):
    """A node filter for nodes with caching enabled"""
    def accept(this, node):
        return node.isCacheEnabled()

cachingnodes = modeler.script.stream().findAll(CacheFilter(), False)
```

例 2: ユーザーの権限に基づき、ディレクトリーまたはファイルの情報をユーザーが取得できるようにする

ユーザーに PSAPI を公開させないようにするために、PSAPI 関数の呼び出しを介して `session.getServerFileSystem()` というメソッドを使用し、ファイル・システム・オブジェクトを作成できます。

以下の例は、IBM スポスモデラー・サーバー に接続するユーザーの権限に基づいて、ディレクトリーまたはファイルの情報をユーザーが取得できるようにする方法を示しています。

```
import modeler.api
stream = modeler.script.stream()
sourceNode = stream.findByID('')
session = modeler.script.session()
fileSystem = session.getServerFileSystem()
parameter = stream.getParameterValue('VPATH')
serverDirectory = fileSystem.getServerFile(parameter)
files = fileSystem.GetFiles(serverDirectory)
```

```

for f in files:
    if f.isDirectory():
        print 'Directory:'
    else:
        print 'File:'
        sourceNode.setPropertyValue('full_filename',f.getPath())
        break
    print f.getName(),f.getPath()
stream.execute()

```

メタデータ: データに関する情報

ストリーム内では複数のノードが互いに接続されているため、各ノードで使用可能な列またはフィールドに関する情報を使用できます。これにより、例えば **Modeler UI** では、ソートまたは集計の基準となるフィールドを選択できます。この情報はデータ・モデルと呼ばれます。

スクリプトは、ノードを出入りするフィールドを調べることによって、データ・モデルにアクセスすることも可能です。一部のノードでは、入力データ・モデルと出力データ・モデルが同じです。例えば、ソート・ノードは、レコードを並べ替えるだけで、データ・モデルを変更することはありません。一部のノード (フィールド作成ノードなど) では、新しいフィールドを追加できます。他のノード (フィルター・ノードなど) は、フィールドの名前を変更したり、フィールドを削除したりすることができます。

以下の例では、スクリプトは標準の **IBM SPSS Modeler druglearn.str** ストリームを使用し、フィールドごとに、いずれかの入力フィールドがドロップされたモデルを作成します。これは、以下のように行われます。

1. データ型ノードから出力データ・モデルにアクセスする。
2. 出力データ・モデルの各フィールドをループする。
3. 各入力フィールドのフィルター・ノードを変更する。
4. 構築中のモデルの名前を変更する。
5. モデル構築ノードを実行する。

注: **druglearn.str** ストリームでスクリプトを実行する前に、スクリプト言語を **Python** に設定することを忘れないでください (ストリームは以前のバージョンの **IBM SPSS Modeler** で作成されているため、ストリーム・スクリプト言語はレガシーに設定されます)。

```

import modeler.api

stream = modeler.script.stream()
filternode = stream.findByType("filter", None)
typenode = stream.findByType("type", None)
c50node = stream.findByType("c50", None)
# Always use a custom model name
c50node.setPropertyValue("use_model_name", True)

lastRemoved = None
fields = typenode.getOutputDataModel()
for field in fields:
    # If this is the target field then ignore it
    if field.getModelingRole() == modeler.api.ModelingRole.OUT:
        continue

    # Re-enable the field that was most recently removed
    if lastRemoved != None:
        filternode.setKeyedPropertyValue("include", lastRemoved, True)

    # Remove the field
    lastRemoved = field.getColumnname()
    filternode.setKeyedPropertyValue("include", lastRemoved, False)

    # Set the name of the new model then run the build
    c50node.setPropertyValue("model_name", "Exclude " + lastRemoved)
    c50node.run([])

```

DataModel オブジェクトには、データ・モデル内のフィールドまたは列に関する情報にアクセスするための多くのメソッドがあります。これらのメソッドを次の表に要約します。

表 18. フィールドまたは列に関する情報にアクセスするための <i>DataModel</i> オブジェクト・メソッド		
方法	戻り値の型	説明
<code>d.getColumnCount()</code>	<i>int</i>	データ・モデル内の列の数を返します。
<code>d.columnIterator()</code>	イテレーター	各列を「ファイル順」の挿入順序で返すイテレーターを返します。イテレーターは列のインスタンスを返します。
<code>d.nameIterator()</code>	イテレーター	各列の名前を「ファイル順」の挿入順序で返すイテレーターを返します。
<code>d.contains(name)</code>	<i>Boolean</i>	指定した名前の列がこの <i>DataModel</i> 内に存在する場合は <i>True</i> を返し、存在しない場合は <i>False</i> を返します。
<code>d.getColumn(name)</code>	列	指定された名前の列を戻します。
<code>d.getColumnGroup(name)</code>	<i>ColumnGroup</i>	指定した列グループを返すか、指定した列グループが存在しない場合は <i>None</i> を返します。
<code>d.getColumnGroupCount()</code>	<i>int</i>	このデータ・モデル内の列グループの数を返します。
<code>d.columnGroupIterator()</code>	イテレーター	各列グループを順番に返すイテレーターを返します。
<code>d.toArray()</code>	列 []	データ・モデルを列の配列として返します。列は「ファイル順」の挿入順序になります。

各フィールド (*Column* オブジェクト) には、列に関する情報にアクセスするための多くのメソッドが含まれています。以下の表に、これらのメソッドを示します。

表 19. 列に関する情報にアクセスするための <i>Column</i> オブジェクト・メソッド		
方法	戻り値の型	説明
<code>c.columnName()</code>	<i>string</i>	列の名前を戻します。
<code>c.columnLabel()</code>	<i>string</i>	列のラベルを返すか、列にラベルが関連付けられていない場合は空文字列を返します。
<code>c.measureType()</code>	<i>MeasureType</i>	列の測定タイプを返します。
<code>c.storageType()</code>	<i>StorageType</i>	列のストレージ・タイプを返します。
<code>c.isMeasureDiscrete()</code>	<i>Boolean</i>	列が離散型の場合は <i>True</i> を返します。セット型またはフラグ型の列は、離散型と見なされます。
<code>c.isModelOutputColumn()</code>	<i>Boolean</i>	列がモデル出力列の場合は <i>True</i> を返します。

表 19. 列に関する情報にアクセスするための *Column* オブジェクト・メソッド (続き)

方法	戻り値の型	説明
<code>c.isStorageDatetime()</code>	<i>Boolean</i>	列のストレージが、時刻、日付、またはタイム・スタンプの値の場合は <i>True</i> を返します。
<code>c.isStorageNumeric()</code>	<i>Boolean</i>	列のストレージが整数または実数の場合は <i>True</i> を返します。
<code>c.isValidValue(value)</code>	<i>Boolean</i>	指定した値がこのストレージで有効な場合は <i>True</i> を返し、有効な列の値が分かる場合は <i>valid</i> を返します。
<code>c.getModelingRole()</code>	<i>ModelingRole</i>	列のモデル作成の役割を返します。
<code>c.getSetValues()</code>	オブジェクト []	列の有効な値の配列を返すか、値が分からない場合または列がセット型でない場合は <i>None</i> を返します。
<code>c.getValueLabel(value)</code>	<i>string</i>	列の値のラベルを返すか、値にラベルが関連付けられていない場合は空文字列を返します。
<code>c.getFalseFlag()</code>	オブジェクト	列の「false」標識値を返すか、値が分からない場合または列がフラグ型でない場合は <i>None</i> を返します。
<code>c.getTrueFlag()</code>	オブジェクト	列の「true」標識値を返すか、値が分からない場合または列がフラグ型でない場合は <i>None</i> を返します。
<code>c.getLowerBound()</code>	オブジェクト	列の値の下限値を返すか、値が分からない場合または列が連続型でない場合は <i>None</i> を返します。
<code>c.getUpperBound()</code>	オブジェクト	列の値の上限値を返すか、値が分からない場合または列が連続型でない場合は <i>None</i> を返します。

列に関する情報にアクセスするほとんどのメソッドには、*DataModel* オブジェクトに定義されている同等のメソッドがあります。たとえば、次の 2 つのステートメントは、同じものを指します。

```
dataModel.getColumn("someName").getModelingRole()
dataModel.getModelingRole("someName")
```

生成されたオブジェクトへのアクセス

ストリームを実行するには、通常、追加の出力オブジェクトを生成する必要があります。これらの追加のオブジェクトは、新規モデル (以降の実行で使用する情報を提供する出力) にすることができます。

下記の例では、ストリームの開始点として `druglearn.str` ストリームを再度使用しています。この例では、ストリームのすべてのノードを実行し、結果をリストに格納します。その後、スクリプトは結果をル

ープし、実行結果のモデル出力は IBM SPSS Modeler モデル (.gm) ファイルとして保存され、モデルは PMML エクスポートされます。

```
import modeler.api

stream = modeler.script.stream()

# Set this to an existing folder on your system.
# Include a trailing directory separator
modelFolder = "C:/temp/models/"

# Execute the stream
models = []
stream.runAll(models)

# Save any models that were created
taskrunner = modeler.script.session().getTaskRunner()
for model in models:
    # If the stream execution built other outputs then ignore them
    if not(isinstance(model, modeler.api.ModelOutput)):
        continue

    label = model.getLabel()
    algorithm = model.getModelDetail().getAlgorithmName()

    # save each model...
    modelFile = modelFolder + label + algorithm + ".gm"
    taskrunner.saveModelToFile(model, modelFile)

    # ...and export each model PMML...
    modelFile = modelFolder + label + algorithm + ".xml"
    taskrunner.exportModelToFile(model, modelFile, modeler.api.FileFormat.XML)
```

タスク実行クラスは、よく使用するさまざまな処理を実行するのに便利です。このクラスで使用可能なメソッドの要約を以下の表に示します。

表 20. よく使用する処理を実行するためのタスク実行クラスのメソッド		
方法	戻り値の型	説明
<code>t.createStream(name, autoConnect, autoManage)</code>	ストリーム	新規ストリームを作成して返します。非公開でストリームを作成してユーザーから不可視にする必要があるコードでは、 <code>autoManage</code> フラグを <code>False</code> に設定する必要があります。
<code>t.exportDocumentToFile(documentOutput, filename, fileFormat)</code>	適用外	指定されたファイル形式を使用してストリームの説明をファイルにエクスポートします。
<code>t.exportModelToFile(modelOutput, filename, fileFormat)</code>	適用外	指定されたファイル形式を使用してモデルをファイルにエクスポートします。
<code>t.exportStreamToFile(stream, filename, fileFormat)</code>	適用外	指定されたファイル形式を使用してストリームをファイルにエクスポートします。
<code>t.insertNodeFromFile(filename, diagram)</code>	ノード	指定されたファイルからノードを読み込み、指定されたダイアグラムに挿入して返します。ノード・オブジェクトとスーパーノード・オブジェクトの両方の読み込みに使用することができます。
<code>t.openDocumentFromFile(filename, autoManage)</code>	DocumentOutput	指定されたファイルからドキュメントを読み込んで返します。

表 20. よく使用する処理を実行するためのタスク実行クラスのメソッド (続き)

方法	戻り値の型	説明
<code>t.openModelFromFile(filename, autoManage)</code>	ModelOutput	指定されたファイルからモデルを読み込んで返します。
<code>t.openStreamFromFile(filename, autoManage)</code>	ストリーム	指定されたファイルからストリームを読み込んで返します。
<code>t.saveDocumentToFile(documentOutput, filename)</code>	適用外	指定されたファイルの場所にドキュメントを保存します。
<code>t.saveModelToFile(modelOutput, filename)</code>	適用外	指定されたファイルの場所にモデルを保存します。
<code>t.saveStreamToFile(stream, filename)</code>	適用外	指定されたファイルの場所にストリームを保存します。

エラーの処理

Python 言語には、`try...except` コード・ブロックによる エラー処理が備わっています。スクリプト内でこれを使用すると、例外をトラップし、対処しなければスクリプトが終了してしまう問題を処理することができます。

下記のスクリプト例では、IBM SPSS Collaboration and Deployment Services Repository からモデルを取得しようとしています。この操作では例外が発生する可能性があります (例えば、リポジトリのログイン資格情報が正しく設定されていない場合や、リポジトリのパスが誤っている場合が考えられます)。このスクリプトでは、`ModelerException` がスローされることがあります (IBM SPSS Modeler によって生成されるすべての例外は `modeler.api.ModelerException` から派生します)。

```
import modeler.api

session = modeler.script.session()
try:
    repo = session.getRepository()
    m = repo.retrieveModel("/some-non-existent-path", None, None, True)
    # print goes to the Modeler UI script panel Debug tab
    print "Everything OK"
except modeler.api.ModelerException, e:
    print "An error occurred:", e.getMessage()
```

注: スクリプト操作によっては、標準の Java 例外が発生する場合があります。それらの例外は `ModelerException` から派生していません。それらの例外をキャッチするために、追加の `except` ブロックを使用してすべての Java 例外をキャッチすることができます。以下に例を示します。

```
import modeler.api
import java.lang.Exception

session = modeler.script.session()
try:
    repo = session.getRepository()
    m = repo.retrieveModel("/some-non-existent-path", None, None, True)
    # print goes to the Modeler UI script panel Debug tab
    print "Everything OK"
except modeler.api.ModelerException, e:
    print "An error occurred:", e.getMessage()
except java.lang.Exception, e:
    print "A Java exception occurred:", e.getMessage()
```

ストリーム、セッション、およびスーパーノード・パラメーター

パラメーターは、直接スクリプトの中で値を固定的にコーディングするのではなく、実行時に渡す場合に便利です。パラメーターとその値は、ストリームの場合と同じ方法で定義します。つまり、ストリームまたはスーパーノードのパラメーター・テーブルの項目として、またはコマンド・ラインのパラメーターと

して定義します。以下の表に示すように、Stream クラスおよび SuperNode クラスは、ParameterProvider オブジェクトによって定義される一連の関数を実装しています。セッションには `getParameters()` の呼び出しが用意されており、呼び出すと、それらの関数を定義するオブジェクトが返されます。

表 21. ParameterProvider オブジェクトによって定義されている関数		
方法	戻り値の型	説明
<code>p.parameterIterator()</code>	イテレーター	このオブジェクトのパラメーター名の反復子を返します。
<code>p.getParameterDefinition(parameterName)</code>	ParameterDefinition	指定された名前を持つパラメーターのパラメーター定義を返します。該当するパラメーターがこのプロバイダーに存在しない場合は <code>None</code> を返します。結果は、メソッドが呼び出された時点での定義のスナップショットである可能性があり、その後このプロバイダーを通じてパラメーターに対して行われた変更が反映されているとは限りません。
<code>p.getParameterLabel(parameterName)</code>	string	指定されたパラメーターのラベルを返します。該当するパラメーターが存在しない場合は <code>None</code> を返します。
<code>p.setParameterLabel(parameterName, label)</code>	適用外	指定されたパラメーターのラベルを設定します。
<code>p.getParameterStorage(parameterName)</code>	ParameterStorage	指定されたパラメーターのストレージを返します。該当するパラメーターが存在しない場合は <code>None</code> を返します。
<code>p.setParameterStorage(parameterName, storage)</code>	適用外	指定されたパラメーターのストレージを設定します。
<code>p.getParameterType(parameterName)</code>	ParameterType	指定されたパラメーターのデータ型を返します。該当するパラメーターが存在しない場合は <code>None</code> を返します。
<code>p.setParameterType(parameterName, type)</code>	適用外	指定されたパラメーターのデータ型を設定します。
<code>p.getParameterValue(parameterName)</code>	オブジェクト	指定されたパラメーターの値を返します。該当するパラメーターが存在しない場合は <code>None</code> を返します。
<code>p.setParameterValue(parameterName, value)</code>	適用外	指定されたパラメーターの値を設定します。

以下の例では、スクリプトで通信データを集計して、平均収入データが最も低い領域を探します。次に、その領域でストリーム・パラメーターを設定します。さらに、そのストリーム・パラメーターを条件抽出ノードで使用してその領域をデータから除外した後、残りのデータに対する顧客離れモデルを作成します。

この例では、スクリプトで条件抽出ノード自体を生成するため、正しい値を条件抽出ノードの式に直接生成できたという点で、不自然な例になっています。しかし、通常ストリームは事前に作成されているため、この方法でパラメーターを設定すると便利です。

スクリプト例の最初の部分では、平均収入が最も低い領域を格納するストリーム・パラメーターを作成します。また、スクリプトでは集計ブランチとモデル作成ブランチにノードを作成し、相互に接続します。

```
import modeler.api

stream = modeler.script.stream()

# Initialize a stream parameter
stream.setParameterStorage("LowestRegion", modeler.api.ParameterStorage.INTEGER)

# First create the aggregation branch to compute the average income per region
statisticsimportnode = stream.createAt("statisticsimport", "SPSS File", 114, 142)
statisticsimportnode.setPropertyValue("full_filename", "$CLEO_DEMOS/telco.sav")
statisticsimportnode.setPropertyValue("use_field_format_for_storage", True)

aggregatenode = modeler.script.stream().createAt("aggregate", "Aggregate", 294, 142)
aggregatenode.setPropertyValue("keys", ["region"])
aggregatenode.setKeyedPropertyValue("aggregates", "income", ["Mean"])

tablenode = modeler.script.stream().createAt("table", "Table", 462, 142)

stream.link(statisticsimportnode, aggregatenode)
stream.link(aggregatenode, tablenode)

selectnode = stream.createAt("select", "Select", 210, 232)
selectnode.setPropertyValue("mode", "Discard")
# Reference the stream parameter in the selection
selectnode.setPropertyValue("condition", "'region' = '$P-LowestRegion'")

typenode = stream.createAt("type", "Type", 366, 232)
typenode.setKeyedPropertyValue("direction", "churn", "Target")

c50node = stream.createAt("c50", "C5.0", 534, 232)

stream.link(statisticsimportnode, selectnode)
stream.link(selectnode, typenode)
stream.link(typenode, c50node)
```

このスクリプト例では以下のストリームを作成します。

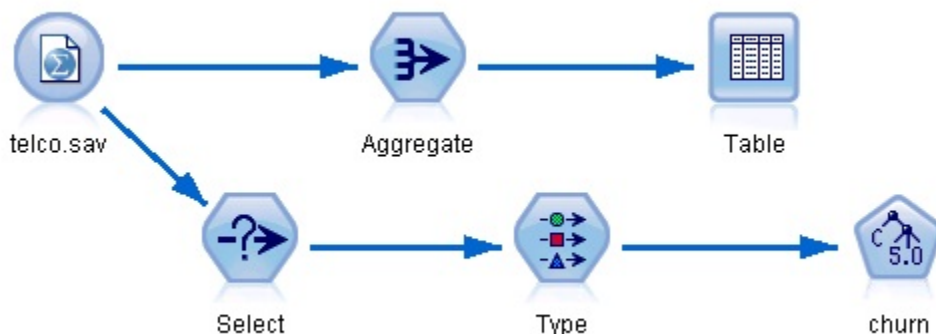


図 5. スクリプト例から得られるストリーム

スクリプト例の以下の部分では、集計ブランチの終端でテーブル・ノードを実行します。

```
# First execute the table node
results = []
tablenode.run(results)
```


スクリプト例の以下の部分では、テーブル・ノードの実行によって生成されたテーブル出力にアクセスします。スクリプトでは次に、テーブルの行全体について反復し、平均収入が最も低い領域を探します。

```
# Running the table node should produce a single table as output
table = results[0]

# table output contains a RowSet so we can access values as rows and columns
rowset = table.getRowSet()
min_income = 10000000.0
min_region = None

# From the way the aggregate node is defined, the first column
# contains the region and the second contains the average income
row = 0
rowcount = rowset.getRowCount()
while row < rowcount:
    if rowset.getValueAt(row, 1) < min_income:
        min_income = rowset.getValueAt(row, 1)
        min_region = rowset.getValueAt(row, 0)
    row += 1
```

スクリプトの以下の部分では、平均収入が最も低い領域を使用して、以前に作成した「LowestRegion」ストリーム・パラメーターを設定します。スクリプトでは次に、指定の領域を学習データから除外してモデル・ビルダーを実行します。

```
# Check that a value was assigned
if min_region != None:
    stream.setParameterValue("LowestRegion", min_region)
else:
    stream.setParameterValue("LowestRegion", -1)

# Finally run the model builder with the selection criteria
c50node.run([])
```

スクリプト例全体を以下に示します。

```
import modeler.api

stream = modeler.script.stream()

# Create a stream parameter
stream.setParameterStorage("LowestRegion", modeler.api.ParameterStorage.INTEGER)

# First create the aggregation branch to compute the average income per region
statisticsimportnode = stream.createAt("statisticsimport", "SPSS File", 114, 142)
statisticsimportnode.setPropertyValue("full_filename", "$CLEO_DEMOS/telco.sav")
statisticsimportnode.setPropertyValue("use_field_format_for_storage", True)

aggregatenode = modeler.script.stream().createAt("aggregate", "Aggregate", 294, 142)
aggregatenode.setPropertyValue("keys", ["region"])
aggregatenode.setKeyedPropertyValue("aggregates", "income", ["Mean"])

tablenode = modeler.script.stream().createAt("table", "Table", 462, 142)

stream.link(statisticsimportnode, aggregatenode)
stream.link(aggregatenode, tablenode)

selectnode = stream.createAt("select", "Select", 210, 232)
selectnode.setPropertyValue("mode", "Discard")
# Reference the stream parameter in the selection
selectnode.setPropertyValue("condition", "'region' = '$P-LowestRegion'")

typenode = stream.createAt("type", "Type", 366, 232)
typenode.setKeyedPropertyValue("direction", "churn", "Target")

c50node = stream.createAt("c50", "C5.0", 534, 232)

stream.link(statisticsimportnode, selectnode)
stream.link(selectnode, typenode)
stream.link(typenode, c50node)

# First execute the table node
results = []
tablenode.run(results)

# Running the table node should produce a single table as output
```

```

table = results[0]

# table output contains a RowSet so we can access values as rows and columns
rowset = table.getRowSet()
min_income = 1000000.0
min_region = None

# From the way the aggregate node is defined, the first column
# contains the region and the second contains the average income
row = 0
rowcount = rowset.getRowCount()
while row < rowcount:
    if rowset.getValueAt(row, 1) < min_income:
        min_income = rowset.getValueAt(row, 1)
        min_region = rowset.getValueAt(row, 0)
    row += 1

# Check that a value was assigned
if min_region != None:
    stream.setParameterValue("LowestRegion", min_region)
else:
    stream.setParameterValue("LowestRegion", -1)

# Finally run the model builder with the selection criteria
c50node.run([])

```

グローバル値

グローバル値は、指定したフィールドの各種の要約統計量を計算するために使用します。これらの要約値には、ストリーム内の任意の場所からアクセスできます。グローバル値は、ストリームから名前でアクセスできるという点でストリーム・パラメーターと似ています。ストリーム・パラメーターとの相違点は、スクリプトやコマンド・ラインから代入するのではなく、グローバルの設定ノードが実行されると関連付けられた値が自動的に更新されることです。ストリームのグローバル値にアクセスするには、ストリームの `getGlobalValues()` メソッドを呼び出します。

GlobalValues オブジェクトは、以下の表に示す関数を定義しています。

表 22. GlobalValues オブジェクトによって定義されている関数		
方法	戻り値の型	説明
<code>g.fieldNameIterator()</code>	イテレーター	グローバル値を 1 つ以上持つ各フィールド名の反復子を返します。
<code>g.getValue(type, fieldName)</code>	オブジェクト	指定されたデータ型およびフィールド名のグローバル値を返します。値が見つからない場合は <code>None</code> を返します。返される値は一般に数値ですが、将来の実装では別の型の値を返すようになる可能性があります。
<code>g.getValues(fieldName)</code>	マップ	指定されたフィールド名の既知のエントリーを含むマップを返します。フィールドに既存のエントリーがない場合は <code>None</code> を返します。

GlobalValues.Type は、使用可能な要約統計量のタイプを定義します。以下の要約統計量が使用可能です。

- MAX: フィールドの最大値。
- MEAN: フィールドの平均値。
- MIN: フィールドの最小値。
- STDDEV: フィールドの標準偏差。

- SUM: フィールドの値の合計。

例えば、以下のスクリプトは「income」フィールドの平均値にアクセスします。このフィールドは、グローバルの設定ノードによって計算されます。

```
import modeler.api

globals = modeler.script.stream().getGlobalValues()
mean_income = globals.getValue(modeler.api.GlobalValues.Type.MEAN, "income")
```

複数のストリームの処理: スタンドアロン スクリプト

複数のストリームを処理するには、スタンドアロン スクリプトを使用する必要があります。スタンドアロン スクリプトは、IBM スポス モデラー UI 内で編集して実行するか、バッチ・モードでコマンド・ライン・パラメーターとして渡すことができます。

以下のスタンドアロン スクリプトは2つのストリームを開きます。一方のストリームはモデルを作成し、2番目のストリームは予測値の分布をプロットします。

```
# Change to the appropriate location for your system
demosDir = "C:/Program Files/IBM/SPSS/Modeler/18.4.0/DEMOS/streams/"

session = modeler.script.session()
tasks = session.getTaskRunner()

# Open the model build stream, locate the C5.0 node and run it
buildstream = tasks.openStreamFromFile(demosDir + "druglearn.str", True)
c50node = buildstream.findByType("c50", None)
results = []
c50node.run(results)

# Now open the plot stream, find the Na_to_K derive and the histogram
plotstream = tasks.openStreamFromFile(demosDir + "drugplot.str", True)
derivenode = plotstream.findByType("derive", None)
histogramnode = plotstream.findByType("histogram", None)

# Create a model applier node, insert it between the derive and histogram nodes
# then run the histogram
applyc50 = plotstream.createModelApplier(results[0], results[0].getName())
applyc50.setPositionBetween(derivenode, histogramnode)
plotstream.linkBetween(applyc50, derivenode, histogramnode)
histogramnode.setPropertyValue("color_field", "$C-Drug")
histogramnode.run([])

# Finally, tidy up the streams
buildstream.close()
plotstream.close()
```

さらに次の例は、開いているストリーム（「ストリーム」タブで開いているすべてのストリーム）を反復処理する方法を示しています。これは、スタンドアロン スクリプトでのみサポートされることに注意してください。

```
for stream in modeler.script.streams():
    print stream.getName()
```


第5章 スクリプトのヒント

このセクションでは、スクリプトのヒントと使い方について概要を説明します。これには、ストリームの実行を修正したり、スクリプトで暗号化されたパスワードを使用したり、また、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・リポジトリでオブジェクトにアクセスしたりする作業が含まれます。

ストリーム実行の変更

ストリームを実行すると、ターミナル・ノードがデフォルトの状況に最適化された順番で実行されます。状況に応じて、別の順序で実行させることもできます。ストリームの実行順序を変更するには、「ストリームのプロパティ」ダイアログ・ボックスの「実行」タブで、以下の手順を実行します。

1. 空のスクリプトを用意します。
2. ツールバーの「**デフォルト スクリプトを追加**」ボタンをクリックして、デフォルトのストリーム・スクリプトを追加します。
3. デフォルトのストリーム・スクリプトの文の順序を、実際に実行する順序に変更します。

ノードのループ

for ループを使用して、ストリーム内のすべてのノードをループできます。例えば、以下のスクリプト例はすべてのノードをループし、フィルター・ノードにおけるフィールド名を大文字に変更します。

除外されるフィールドが何もなくても、このスクリプトはフィルター・ノードを持つどのようなストリームにおいても使用できます。フィールド名を全面的に大文字に変更するには、すべてのフィールドを渡すフィルター・ノードをただ単に追加するだけです。

```
# Alternative 1: using the data model nameIterator() function
stream = modeler.script.stream()
for node in stream.iterator():
    if (node.getTypeName() == "filter"):
        # nameIterator() returns the field names
        for field in node.getInputDataModel().nameIterator():
            newname = field.upper()
            node.setKeyedPropertyValue("new_name", field, newname)

# Alternative 2: using the data model iterator() function
stream = modeler.script.stream()
for node in stream.iterator():
    if (node.getTypeName() == "filter"):
        # iterator() returns the field objects so we need
        # to call getColumnName() to get the name
        for field in node.getInputDataModel().iterator():
            newname = field.getColumnName().upper()
            node.setKeyedPropertyValue("new_name", field.getColumnName(), newname)
```

このスクリプトは現在のストリーム内のすべてのノードをループし、各ノードがフィルターであるかどうかをチェックします。ノードがフィルターである場合、そのノードの各フィールドをループし、`field.upper()` 関数または `field.getColumnName().upper()` 関数を使用して、名前を大文字に変更します。

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・リポジトリ内のオブジェクトへのアクセス

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・リポジトリのライセンスを保有している場合は、スクリプト コマンドを使用して、オブジェクトをリポジトリに格納したり、リポジトリから取得したりできます。リポジトリを使用して、エンタープライズアプリケーション、ツール、およびソリュ

ーションのコンテキストで、データ マイニング モデルおよび関連する予測オブジェクトのライフサイクルを管理します。

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・リポジトリ への接続

リポジトリにアクセスするには、まず、スポス モデラー ユーザー インターフェースの「ツール」メニューまたはコマンドラインから、リポジトリに対して有効な接続を設定する必要があります。詳しくは、[69 ページの『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 接続の引数』](#)を参照してください。

リポジトリへのアクセスの実行

リポジトリにはセッションからアクセスできます。例えば、以下のようにします。

```
repo = modeler.script.session().getRepository()
```

リポジトリからのオブジェクトの取得

スクリプト内で、ストリーム、モデル、出力、およびノードなど、さまざまなオブジェクトにアクセスするには、`retrieve*` 関数を使用します。取得関数の要約を以下の表に示します。

表 23. 取得スクリプト関数	
オブジェクト・タイプ	リポジトリ関数
ストリーム	<code>repo.retrieveStream(String path, String version, String label, Boolean autoManage)</code>
モデル	<code>repo.retrieveModel(String path, String version, String label, Boolean autoManage)</code>
出力	<code>repo.retrieveDocument(String path, String version, String label, Boolean autoManage)</code>
ノード	<code>repo.retrieveProcessor(String path, String version, String label, ProcessorDiagram diagram)</code>

例えば、以下の関数を使用してリポジトリからストリームを取得できます。

```
stream = repo.retrieveStream("/projects/retention/risk_score.str", None, "production", True)
```

この例は、指定したフォルダーから `risk_score.str` ストリームを取得します。ラベル `production` は、取得するストリームのバージョンを識別します。最後のパラメーターは、スポス モデラー がストリームを管理することを指定します (例えば、スポス モデラー ユーザー・インターフェースが表示されている場合は、**ストリーム** タブにストリームが表示されます)。代わりに、特定の、ラベル付けのないバージョンを使用するには、以下のようにします。

```
stream = repo.retrieveStream("/projects/retention/risk_score.str", "0:2015-10-12 14:15:41.281", None, True)
```

注：バージョンとラベルの両方のパラメーターが `None` の場合、最新バージョンが返されます。

リポジトリへのオブジェクトの格納

スクリプトを使用してリポジトリにオブジェクトを格納するには、`store*` 関数を使用します。格納関数の要約を以下の表に示します。

表 24. 格納スクリプト関数	
オブジェクト・タイプ	リポジトリ関数
ストリーム	<code>repo.storeStream(ProcessorStream stream, String path, String label)</code>
モデル	<code>repo.storeModel(ModelOutput modelOutput, String path, String label)</code>
出力	<code>repo.storeDocument(DocumentOutput documentOutput, String path, String label)</code>
ノード	<code>repo.storeProcessor(Processor node, String path, String label)</code>

例えば、以下の関数を使用して `risk_score.str` ストリームの新規バージョンを格納できます。

```
versionId = repo.storeStream(stream, "/projects/retention/risk_score.str", "test")
```

この例は、新規バージョンのストリームを格納して、それに `"test"` ラベルを関連付けて、新規に作成されたバージョンのバージョン マーカーを返します。

注: ラベルを新規バージョンと関連付けたくない場合は、ラベルには `None` を渡してください。

リポジトリ フォルダの管理

リポジトリ内でフォルダーを使用することで、オブジェクトを論理グループ別に整理でき、オブジェクトの関連がわかりやすくなります。以下の例のように、`createFolder()` 関数を使用してフォルダーを作成してください。

```
newpath = repo.createFolder("/projects", "cross-sell")
```

この例では、`"/projects"` フォルダー内に `"cross-sell"` という名前の新規フォルダーを作成します。この関数は、新規フォルダーの完全パスを返します。

フォルダーの名前を変更するには、以下のように `renameFolder()` 関数を使用してください。

```
repo.renameFolder("/projects/cross-sell", "cross-sell-Q1")
```

最初のパラメーターは名前変更されるフォルダーの完全パスであり、2 番目のパラメーターはそのフォルダーにつける新しい名前です。

空のフォルダーを削除するには、以下のように `deleteFolder()` 関数を使用してください。

```
repo.deleteFolder("/projects/cross-sell")
```

オブジェクトのロックおよびロック解除

スクリプトから、オブジェクトをロックして、ほかのユーザーが既存のバージョンを更新したり新しいバージョンを作成しないようにすることができます。ロックされたオブジェクトのロックを解除することもできます。

オブジェクトをロックおよびロック解除するシンタックスは次のとおりです。

```
repo.lockFile(REPOSITORY_PATH)
repo.lockFile(URI)

repo.unlockFile(REPOSITORY_PATH)
repo.unlockFile(URI)
```

オブジェクトの保存および取得同様、REPOSITORY_PATHによって、リポジトリ内のオブジェクトの場所が決まります。パスを引用符で囲み、区切り文字としてスラッシュを使用する必要があります。大文字と小文字は区別しません。

```
repo.lockFile("/myfolder/Stream1.str")
repo.unlockFile("/myfolder/Stream1.str")
```

また、オブジェクトの場所を決めるには、リポジトリ・パスではなく URI (Uniform Resource Identifier) を使用できます。URI は接頭辞 `spsscr:` を含み、完全に引用符で囲まれている必要があります。パス区切り文字としてはスラッシュだけを使うことができ、スペースは暗号化する必要があります。つまり、パス内ではスペースの代わりに `%20` を使用します。URI では、大文字と小文字は区別しません。いくつか例を挙げると次の通りです。

```
repo.lockFile("spsscr:///myfolder/Stream1.str")
repo.unlockFile("spsscr:///myfolder/Stream1.str")
```

オブジェクトのロックはすべてのバージョンのオブジェクトに適用されます。各バージョンをロックまたはロック解除することはできません。

暗号化パスワードの生成

場合によっては、スクリプトにパスワードを記述する必要があるかもしれません。例えば、パスワードで保護されたデータ・ソースにアクセスしたい場合などです。暗号化パスワードは、次の場所で使用することができます。

- データベース入力ノードおよび出力ノードのノード・プロパティ。
- サーバーにログインするためのコマンド・ライン引数。
- エクスポート・ノードの「公開」タブから生成するパラメーター・ファイル `.par` ファイルに保管されるデータベース接続プロパティ。

ユーザー・インターフェースから、Blowfish アルゴリズムに基づいた暗号化パスワードを生成することができます (詳細については、<http://www.schneier.com/blowfish.html> を参照してください)。パスワードを暗号化したら、そのパスワードをコピーしてスクリプト・ファイルやコマンド・ライン引数に指定することができます。 `databasenode` および `databaseexportnode` に使用するノード・プロパティ `epassword` は暗号化パスワードを格納します。

1. 暗号化パスワードを生成するには、「ツール」メニューから次の項目を選択します。
「パスワードのエンコード...」
2. 「パスワード」ボックスにパスワードを指定します。
3. 「暗号化」をクリックすると、ランダムに暗号化されたパスワードが生成されます。
4. 「コピー」ボタンをクリックすると、暗号化されたパスワードがクリップボードにコピーされます。
5. パスワードを目的のスクリプトやパラメーターに貼り付けます。

スクリプトの検査

「スタンドアロン スクリプト」ダイアログ・ボックスのツールバーにある赤い検査ボタンをクリックすれば、すべてのスクリプトのシンタックスを検査することができます。



図 6. ストリーム・スクリプトのツールバー・アイコン

スクリプトの検査時にコードにエラーがあった場合、エラーを警告するメッセージと推奨する修正方法が表示されます。エラーのある行を表示するには、ダイアログ・ボックスの下部にあるフィードバック情報をクリックしてください。エラーが赤で強調表示されます。

コマンド・ラインからのスクリプト

通常はユーザー・インターフェースから行われるような操作を、スクリプトで実行することができます。IBM スポス モデラー を起動するときには、コマンド・ライン上でスタンドアロン・スクリプトを指定して実行してください。以下に例を示します。

```
client -script scores.txt -execute
```

-script フラグは指定されたスクリプトをロードすることを、-execute フラグはスクリプト・ファイル中のすべてのコマンドを実行することを示しています。

旧リリースとの互換性

以前の IBM スポス モデラー のリリースで作成されたスクリプトは、通常現在のリリースでも変更なしで動作します。ただし、モデル・ナゲットがストリームに自動的に挿入され(デフォルト設定)、ストリーム内のその種類の既存ナゲットを置き換えまたは補足する場合があります。これが実際に起こるかどうかは、**ストリームへのモデルの追加 および 前のモデルを置換 オプション**(「ツール」>「オプション」>「ユーザー・オプション」>「通知」)の設定によって異なります。例えば、既存のナゲットを削除して新しいナゲットを挿入し、ナゲットの置換を処理する旧リリースからのスクリプトの変更が必要な場合があります。

現在のリリースで作成したスクリプトは、以前のリリースでは動作しないことがあります。

古いリリースで作成されたスクリプトがあるコマンドを使用し、そのコマンドがリリースされてから他のコマンドに置き換えられて(または、廃止されて)いる場合は、古い形が依然としてサポートされますが、同時に警告メッセージも表示されます。例えば、古い `generated` キーワードは `model` に、`clear generated` は `clear generated palette` に置き換えられます。古い形を使うスクリプトは依然として動作しますが、警告も表示されます。

ストリーム実行結果へのアクセス

多くの IBM スポス モデラー ノードで、モデル、グラフ、およびテーブル形式データなどの出力オブジェクトが生成されます。これらの出力の多くには、後続の実行をガイドするためにスクリプトで利用できる有用な値が含まれています。これらの値は、コンテンツ コンテナ(単にコンテナと呼ばれる)にグループ化されます。コンテナには、各コンテナを識別するタグまたは ID を使用してアクセスできます。これらの値にアクセスする方法は、そのコンテナが使用する形式(「コンテンツ モデル」)によって異なります。

例えば、多くの予測モデル出力では、PMML という XML の一種を使用して、モデルに関する情報(各分割でディジション ツリーが使用するフィールドや、ニューラル ネットワーク内のニューロンの接続方法とその強度など)を表現します。PMML を使用するモデル出力では、その情報にアクセスするために使用できる XML コンテンツ モデルを提供します。以下に例を示します。

```
stream = modeler.script.stream()
# Assume the stream contains a single C5.0 model builder node
# and that the datasource, predictors and targets have already been
# set up
modelbuilder = stream.findByType("c50", None)
results = []
modelbuilder.run(results)
modeloutput = results[0]

# check how many contents are there and what are their names
tags = modeloutput.getContentModelTags()
print "Content Model Tags :", tags

# Now that we have the C5.0 model output object, access the
# relevant content model
cm = modeloutput.getContentModel("PMML")
if (cm != None) :
    # The PMML content model is a generic XML-based content model that
    # uses XPath syntax. Use that to find the names of the data fields.
    # The call returns a list of strings match the XPath values
```

```
dataFieldNames = cm.getStringValues("/PMML/DataDictionary/DataField",
"name")
print "Data Field Names:" , dataFieldNames
```

IBM スポス モデラー は、スクリプトで以下のコンテンツ モデルをサポートします。

- **テーブル コンテンツ モデル:** 行と列として表現される単純なテーブル形式データにアクセスできます。
- **XML コンテンツ モデル:** XML 形式で保管されたコンテンツにアクセスできます。
- **JSON コンテンツ モデル:** JSON 形式で保管されたコンテンツにアクセスできます。
- **列統計コンテンツ モデル:** 特定のフィールドに関する統計の要約にアクセスできます。
- **ペアごとの列統計コンテンツ モデル:** 2 つのフィールドの間の統計の要約または 2 つの個別のフィールドの間にある値にアクセスできます。

次のノードには、これらのコンテンツ モデルが含まれていないことに注意してください。

- 時系列
- 判別分析
- SLRM
- TCM
- すべての Python ノード
- すべての Spark ノード
- すべてのデータベース モデル作成ノード
- 拡張モデル
- STP

テーブル コンテンツ モデル

テーブル コンテンツ モデルは、単純な行と列のデータにアクセスするための単純なモデルを提供します。特定の列内の値は、すべてストレージのタイプが同じでなければなりません (例えば、文字列または整数)。

API

表 25. API		
戻る	方法	説明
int	getRowCount()	このテーブル内の行の数を返します。
int	getColumnCount()	このテーブル内の列の数を返します。
String	getColumnName(int columnIndex)	指定された列インデックス位置にある列の名前を返します。列のインデックスは 0 から始まります。
StorageType	getStorageType(int columnIndex)	指定されたインデックス位置にある列のストレージ タイプを返します。列のインデックスは 0 から始まります。
Object	getValueAt(int rowIndex, int columnIndex)	指定された行インデックスおよび列インデックスの位置にある値を返します。行と列のインデックスは 0 から始まります。

表 25. API (続き)		
戻る	方法	説明
void	reset()	このコンテンツ モデルに関連付けられた内部ストレージをすべて消去します。

ノードおよび出力

この表では、このタイプのコンテンツ モデルを含む出力を作成するノードをリストします。

表 26. ノードおよび出力		
ノード名	出力名	コンテナ ID
table	table	"table"

スクリプトの例

```
stream = modeler.script.stream()
from modeler.api import StorageType

# Set up the variable file import node
varfilenode = stream.createAt("variablefile", "DRUG Data", 96, 96)
varfilenode.setPropertyValue("full_filename", "$CLEO_DEMOS/DRUG1n")

# Next create the aggregate node and connect it to the variable file node
aggregatenode = stream.createAt("aggregate", "Aggregate", 192, 96)
stream.link(varfilenode, aggregetenode)

# Configure the aggregate node
aggregatenode.setPropertyValue("keys", ["Drug"])
aggregatenode.setKeyedPropertyValue("aggregates", "Age", ["Min", "Max"])
aggregatenode.setKeyedPropertyValue("aggregates", "Na", ["Mean", "SDev"])

# Then create the table output node and connect it to the aggregate node
tablenode = stream.createAt("table", "Table", 288, 96)
stream.link(aggregatenode, tablenode)

# Execute the table node and capture the resulting table output object
results = []
tablenode.run(results)
tableoutput = results[0]

# Access the table output's content model
tablecontent = tableoutput.getContentModel("table")

# For each column, print column name, type and the first row
# of values from the table content
col = 0
while col < tablecontent.getColumnCount():
    print tablecontent.getColumnName(col), \
          tablecontent.getStorageType(col), \
          tablecontent.getValueAt(0, col)
    col = col + 1
```

スクリプトの「デバッグ」タブには、以下のような出力が表示されます。

```
Age_Min Integer 15
Age_Max Integer 74
Na_Mean Real 0.730851098901
Na_SDev Real 0.116669731242
```

XML コンテンツ モデル

XML コンテンツ モデルでは、XML ベースのコンテンツにアクセスできます。

XML コンテンツ モデルは、XPath 式に基づくコンポーネントにアクセスする機能をサポートします。XPath 式は、呼び出し元がどの要素または属性を必要とするかを定義する文字列です。XML コンテンツ モデルは、さまざまなオブジェクトの作成と、XPath のサポートで通常必要となる式のコンパイルについて、詳細な内容を隠します。これにより、Python スクリプトからの呼び出しが単純になります。

XML コンテンツ モデルには、XML 文書を文字列として返す関数が含まれています。これにより、Python スクリプト ユーザーは、自分にとって望ましい Python ライブラリを使用して XML を解析できます。

API

表 27. API		
戻る	方法	説明
String	getXMLAsString()	XML を文字列として返します。
number	getNumericValue(String xpath)	パスを評価した結果を数値として返します (例えば、パス式に一致する要素の数をカウントします)。
boolean	getBooleanValue(String xpath)	指定されたパス式を評価した結果をブール値として返します。
String	getStringValue(String xpath, String attribute)	指定されたパスに一致する、属性値または XML ノード値のいずれかを返します。
List of strings	getStringValues(String xpath, String attribute)	指定されたパスに一致するすべての属性値または XML ノード値のリストを返します。
List of lists of strings	getValuesList(String xpath, <List of strings> attributes, boolean includeValue)	指定されたパスに一致するすべての属性値のリストを、必要な場合は XML ノード値と共に返します。
Hash table (key:string, value:list of string)	getValuesMap(String xpath, String keyAttribute, <List of strings> attributes, boolean includeValue)	キー属性または XML ノード値をキーとして使用するハッシュ テーブルを返し、また、指定された属性値のリストをテーブル値として返します。
boolean	isNamespaceAware()	XML パーサーが名前空間を認識している必要があるかどうかを返します。デフォルトは False です。
void	setNamespaceAware(boolean value)	XML パーサーが名前空間を認識している必要があるかどうかを設定します。このメソッドでは、後続の呼び出しで変更内容が取得されるようにするために <code>reset()</code> も呼び出します。

表 27. API (続き)		
戻る	方法	説明
void	reset()	このコンテンツ モデルに関連付けられた内部ストレージをすべて消去します (キャッシュされた DOM オブジェクトなど)。

ノードおよび出力

この表では、このタイプのコンテンツ モデルを含む出力を作成するノードをリストします。

表 28. ノードおよび出力		
ノード名	出力名	コンテナ ID
Most model builders	Most generated models	"PMML"
"autodataprep"	n/a	"PMML"

スクリプトの例

コンテンツにアクセスするための Python スクリプトのコードは、以下のようになります。

```
results = []
modelbuilder.run(results)
modeloutput = results[0]
cm = modeloutput.getContentModel("PMML")

dataFieldNames = cm.getStringValues("/PMML/DataDictionary/DataField", "name")
predictedNames = cm.getStringValues("//MiningSchema/
MiningField[@usageType='predicted']", "name")
```

JSON コンテンツ モデル

JSON コンテンツ モデルは、JSON 形式のコンテンツのサポートを提供するために使用されます。このモデルでは、どの値にアクセスするかを呼び出し元が認識していることを前提として、呼び出し元が値を抽出できるようにする基本的な API が提供されます。

API

表 29. API		
戻る	方法	説明
String	getJSONAsString()	JSON コンテンツを文字列として返します。
Object	getObjectAt(<List of cbjecta> path, JSONArtifact artifact) throws Exception	指定されたパスのオブジェクトを返します。指定されたルート成果物がヌルである可能性があり、その場合はコンテンツのルートが使用されます。返される値は、リテラル文字列、整数、実数、またはブール値であるか、あるいは JSON 成果物 (JSON オブジェクトまたは JSON 配列のいずれか) である可能性もあります。

表 29. API (続き)		
戻る	方法	説明
Hash table (key:object, value:object)	getChildValuesAt(<List of object> path, JSONArtifact artifact) throws Exception	パスが JSON オブジェクトを指す場合は、指定されたパスの子値を返します。それ以外の場合はヌルを返します。テーブル内のキーは文字列ですが、関連付けられている値は、リテラル文字列、整数、実数、またはブール値であるか、あるいは JSON 成果物 (JSON オブジェクトまたは JSON 配列のいずれか) である可能性もあります。
List of objects	getChildrenAt(<List of object> path path, JSONArtifact artifact) throws Exception	パスが JSON 配列を指す場合は、指定されたパスのオブジェクトのリストを返します。それ以外の場合はヌルを返します。返される値は、リテラル文字列、整数、実数、またはブール値であるか、あるいは JSON 成果物 (JSON オブジェクトまたは JSON 配列のいずれか) である可能性もあります。
void	reset()	このコンテンツ モデルに関連付けられた内部ストレージをすべて消去します (キャッシュされた DOM オブジェクトなど)。

ノードおよび出力

この表では、このタイプのコンテンツ モデルを含む出力を作成するノードをリストします。

表 30. ノードおよび出力		
ノード名	出力名	コンテナ ID
"applykmeansas"	n/a	"JSON_MV"
"applyxgboosttree"	n/a	"JSON_MV"

サンプル・スクリプト

以下のスクリプトは JSON ファイルを取得します。

```
applykmeansas = stream.findByType("applykmeansas",None)
mvjson = applykmeansas.getContentModel("JSON_MV")
print(mvjson.getJSONAsString())
```

```
applyxgboosttree = stream.findByType("applyxgboosttree",None)
mvjson = applyxgboosttree.getContentModel("JSON_MV")
print(mvjson.getJSONAsString())
```

列統計コンテンツ モデルおよびペアごとの統計コンテンツ モデル

列統計コンテンツ モデルでは、フィールドごとに計算できる統計 (1 変量の統計) にアクセスできます。ペアごとの統計コンテンツ モデルでは、フィールドのペア間で計算できる統計またはフィールド内の値にアクセスできます。

統計の尺度には以下のものがあります。

- Count
- UniqueCount
- ValidCount
- Mean
- Sum
- Min
- Max
- Range
- Variance
- StandardDeviation
- StandardErrorOfMean
- Skewness
- SkewnessStandardError
- Kurtosis
- KurtosisStandardError
- Median
- Mode
- Pearson
- Covariance
- TTest
- FTest

一部の値は単一の列統計の場合のみに該当し、その他の値はペアごとの統計の場合のみに該当します。

これらを生成するノードを以下に示します。

- **記述統計ノード**: 列統計を生成し、相関フィールドが指定されている場合はペアごとの統計を生成できます。
- **データ検査ノード**: 列を生成し、オーバーレイ フィールドが指定されている場合はペアごとの統計を生成できます。
- **平均値ノード**: フィールドのペアを比較するとき、またはあるフィールドの値を他のフィールド要約と比較するときに、ペアごとの統計を生成します。

使用可能なコンテンツ モデルと統計は、その特定のノードの機能とそのノード内の設定の両方によって決まります。

ColumnStatsContentModel API

表 31. ColumnStatsContentModel API		
戻る	方法	説明
List<StatisticType>	getAvailableStatistics()	このモデルで使用可能な統計を返します。必ずしもすべてのフィールドがすべての統計の値を持つわけではありません。
List<String>	getAvailableColumns()	統計が計算された対象の列名を返します。
Number	getStatistic(String column, StatisticType statistic)	列に関連付けられた統計値を返します。
void	reset()	このコンテンツ モデルに関連付けられた内部ストレージをすべて消去します。

PairwiseStatsContentModel API

表 32. PairwiseStatsContentModel API		
戻る	方法	説明
List<StatisticType>	getAvailableStatistics()	このモデルで使用可能な統計を返します。必ずしもすべてのフィールドがすべての統計の値を持つわけではありません。
List<String>	getAvailablePrimaryColumns()	統計が計算された対象の 1 次列名を返します。
List<Object>	getAvailablePrimaryValues()	統計が計算された対象の 1 次列の値を返します。
List<String>	getAvailableSecondaryColumns()	統計が計算された対象の 2 次列名を返します。
Number	getStatistic(String primaryColumn, String secondaryColumn, StatisticType statistic)	列に関連付けられた統計値を返します。
Number	getStatistic(String primaryColumn, Object primaryValue, String secondaryColumn, StatisticType statistic)	1 次列値と 2 次列に関連付けられた統計値を返します。
void	reset()	このコンテンツ モデルに関連付けられた内部ストレージをすべて消去します。

ノードおよび出力

この表では、このタイプのコンテンツ モデルを含む出力を作成するノードをリストします。

表 33. ノードおよび出力			
ノード名	出力名	コンテナ ID	注記
"means" (平均値ノード)	"means"	"columnStatistics"	
"means" (平均値ノード)	"means"	"pairwiseStatistics"	
"dataaudit" (データ検査ノード)	"means"	"columnStatistics"	
"statistics" (記述統計ノード)	"statistics"	"columnStatistics"	特定のフィールドが検証された場合のみ生成されます。
"statistics" (記述統計ノード)	"statistics"	"pairwiseStatistics"	フィールドが相関している場合のみ生成されます。

スクリプトの例

```

from modeler.api import StatisticType
stream = modeler.script.stream()

# Set up the input data
varfile = stream.createAt("variablefile", "File", 96, 96)
varfile.setPropertyValue("full_filename", "$CLEO/DEMOS/DRUG1n")

# Now create the statistics node. This can produce both
# column statistics and pairwise statistics
statisticsnode = stream.createAt("statistics", "Stats", 192, 96)
statisticsnode.setPropertyValue("examine", ["Age", "Na", "K"])
statisticsnode.setPropertyValue("correlate", ["Age", "Na", "K"])
stream.link(varfile, statisticsnode)

results = []
statisticsnode.run(results)
statsoutput = results[0]
statscm = statsoutput.getContentModel("columnStatistics")
if (statscm != None):
    cols = statscm.getAvailableColumns()
    stats = statscm.getAvailableStatistics()
    print "Column stats:", cols[0], str(stats[0]), " = ",
    statscm.getStatistic(cols[0], stats[0])

statscm = statsoutput.getContentModel("pairwiseStatistics")
if (statscm != None):
    pcols = statscm.getAvailablePrimaryColumns()
    scols = statscm.getAvailableSecondaryColumns()
    stats = statscm.getAvailableStatistics()
    corr = statscm.getStatistic(pcols[0], scols[0], StatisticType.Pearson)
    print "Pairwise stats:", pcols[0], scols[0], " Pearson = ", corr

```


第 6 章 コマンド・ライン引数

ソフトウェアの起動

オペレーティング・システムのコマンド・ラインを使用し、次のようにして IBM スポス モデラー を起動できます。

1. IBM スポス モデラー がインストールされているコンピューターで、DOS つまりコマンド・プロンプト・ウィンドウを開きます。
2. IBM スポス モデラー インターフェースを対話モードで起動するには、`modelerclient` コマンドに続けて必要な引数を入力します。以下に例を示します。

```
modelerclient -stream report.str -execute
```

使用可能な引数 (フラグ) により、サーバーへの接続、ストリームのロード、スクリプトの実行、または必要に応じて他のパラメーターの指定を行うことができます。

コマンド・ライン引数の使用

コマンド行引数 (フラグとも呼ばれる) を初期 `modelerclient` コマンドに追加して、IBM スポス モデラー の呼び出しを変更することができます。

複数の種類のコマンド・ライン引数を使用できます。これらのコマンド・ライン引数についてはこのセクションで後述します。

表 34. コマンド・ライン引数の種類	
引数の種類	参照箇所
システムの引数	詳しくは、トピック 66 ページの『 システムの引数 』を参照してください。
パラメーターの引数	詳しくは、トピック 67 ページの『 パラメーターの引数 』を参照してください。
サーバー接続の引数	詳しくは、トピック 68 ページの『 サーバー接続の引数 』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・リポジトリ 接続の引数	詳しくは、トピック 69 ページの『 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 接続の引数 』を参照してください。
IBMSPPS Analytic Server 接続の引数	詳しくは、トピック 70 ページの『 IBMSPPS Analytic Server 接続の引数 』を参照してください。

例えば、以下のようにして `-server`、`-stream` および `-execute` のフラグ型を使用してサーバーに接続し、ストリームをロードおよび実行できます。

```
modelerclient -server -hostname myserver -port 80 -username dminer  
-password 1234 -stream mystream.str -execute
```

ローカル・クライアントのインストールと競合する場合、サーバー接続の引数は不要です。

スペースを含むパラメーター値は二重引用符で囲むことができます。例えば、次のようになります。

```
modelerclient -stream mystream.str -Pusername="Joe User" -execute
```

-state フラグと-script フラグをそれぞれ使用して、この方法で IBM スポス モデラー状態とスクリプトを実行することもできます。

注：コマンドで構造化パラメータを使用する場合は、引用符の前に円記号を置く必要があります。これにより、文字列の解釈中に引用符が削除されなくなります。

コマンド・ライン引数のデバッグ

コマンド・ラインをデバッグするには、modelerclient コマンドを使用して、必要な引数を指定して IBM スポス モデラーを起動します。これにより、コマンドが予定通りに実行されることを検証できます。また、「セッションパラメーター」ダイアログ・ボックス（「ツール」メニュー、セッションパラメーターの設定）のコマンド・ラインから渡されるパラメーターの値を確認することもできます。

システムの引数

ユーザー・インターフェースのコマンド・ラインによる起動で利用できるシステム引数を次の表に示します。

表 35. システムの引数	
引数	動作説明
@ <commandFile>	@ 文字に続けてファイル名を記述することにより、コマンド・リストを指定することができます。modelerclient コマンドに @ から始まる引数を指定すると、その引数に指定されたコマンド・ファイル中のコマンドが、コマンド・ラインに指定されているのと同じように処理されます。詳しくは、トピック 70 ページの『複数の引数の組み合わせ』を参照してください。
-directory <dir>	デフォルトの作業ディレクトリを設定します。ローカル・モードでは、このディレクトリはデータと出力の両方で使用されます。例えば、-directory c:/または-directory c:\\
-server_directory <dir>	デフォルトのデータ用サーバー・ディレクトリを設定します。-directory フラグで指定された作業ディレクトリは、出力に使用されます。
-execute	起動後に、起動時にロードされたストリーム、ステート、またはスクリプトを実行します。ストリームやステートではなくスクリプトがロードされた場合は、スクリプトだけが実行されます。
-stream <stream>	起動時に、指定したストリームをロードします。複数のストリームを指定できますが、最後に指定したストリームが現在のストリームに設定されます。
-script <script>	起動時に、指定したスタンドアロン スクリプトをロードします。下で説明しているストリームやステートに加えてこれも指定できますが、起動時には1つのスクリプトしかロードできません。
-model <model>	起動時に、指定の生成モデル(.gm 形式ファイル)をロードします。
-state <state>	起動時に、指定した保存済みのステートをロードします。
-project <project>	指定したプロジェクトをロードします。起動時には、プロジェクトを1つしかロードできません。
-output <output>	起動時に、保存された出力オブジェクト(.cou 形式ファイル)をロードします。
-help	コマンド・ライン引数のリストを表示します。このオプションを指定すると、他の引数はすべて無視されて、ヘルプ画面が表示されます。
-P <名>=<value>	スタートアップ・パラメーターの設定に使用されます。ノードのプロパティ（スロット・パラメーター）の設定に使用することもできます。

注: ユーザー・インターフェースでデフォルト・ディレクトリーも設定できます。このオプションにアクセスするには、「ファイル」メニューの「作業ディレクトリーの設定」または「サーバー ディレクトリーの設定」を選択します。

複数ファイルのロード

ロードされた各オブジェクトに対応する引数を繰り返し指定して、起動時にコマンド・ラインから、複数のストリーム、ステート、および出力をロードすることができます。例えば、`report.str` と `train.str` の2種類のストリームをロード、実行するには、コマンド・ラインに次のコマンドを指定します。

```
modelerclient -stream report.str -stream train.str -execute
```

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・リポジトリーからのオブジェクトのロード

ファイルまたは 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・リポジトリー (ライセンスがある場合) から特定のオブジェクトを読み込むことができるため、ファイル名の接頭辞 `spsscr:` および、オプションで `file:` (ディスク上のオブジェクト) が IBM スポス モデラー にオブジェクトの検索場所を示します。上記の接頭辞は、次のフラグに適用できます。

- `-stream`
- `-script`
- `-output`
- `-model`
- `-project`

接頭辞を使用して、オブジェクトの場所を指定する URI を作成します。例えば、次のようになります。 `-stream "spsscr:///folder_1/scoring_stream.str"`。 `spsscr:` 接頭部が存在するためには、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・リポジトリーへの有効な接続が同じコマンドで指定されている必要があります。そのため、例えば、フル・コマンドは次のようになります。

```
modelerclient -spsscr_hostname myhost -spsscr_port 8080  
-spsscr_username myusername -spsscr_password mypassword  
-stream "spsscr:///folder_1/scoring_stream.str" -execute
```

コマンド・ラインから URI を使用する必要があることに注意してください。より単純な `REPOSITORY_PATH` はサポートされていません。(これはスクリプト内でのみ機能します。) 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・リポジトリー中のオブジェクトの URI 詳細については、[51 ページの『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・リポジトリー内のオブジェクトへのアクセス』](#)を参照してください。

パラメーターの引数

IBM スポス モデラー のコマンド・ライン実行時に、パラメーターをフラグとして使用することができます。コマンド・ライン引数では、`-P <名前>=<値>` の形式のパラメーターを示すために `-P` フラグが使用されます。

パラメーターは、次のいずれかになります。

- **単純なパラメーター** (または、CLEM 式で直接使用されるパラメーター)。
- **スロット・パラメーター** (ノード・プロパティーとも呼ばれる)。これらのパラメーターは、ストリーム中のノードの設定を変更するために使用されます。詳しくは、[75 ページの『ノードのプロパティーの概要』](#)のトピックを参照してください。
- IBM スポス モデラー の起動を変更するために用いられる、**コマンド・ライン・パラメーター**。

例えば、データ・ソースのユーザー名とパスワードを、次のようにコマンド・ラインのフラグとして指定することができます。

```
modelerclient -stream response.str -P:databasenode.datasource="{\"ORA 10gR2\", user1, mypsw, false}"
```

形式は、databasenode ノード・プロパティの datasource パラメーターの形式と同じです。詳しくは、93 ページの『[databasenode プロパティ](#)』を参照してください。

エンコードしたパスワードを渡す場合は、最後のパラメーターを true に設定しなければなりません。また、データベースのユーザー名とパスワードの前にはスペースを入れないでください (ただし、ユーザー名やパスワードに実際に先行スペースが含まれる場合は、この限りではありません)。

注: ノードの名前を指定する場合、二重引用符でノード名を囲み、それらの引用符を円記号でエスケープする必要があります。例えば、直前の例のデータ入力ノード名が Source_ABC である場合、入力は以下のようになります。

```
modelerclient -stream response.str  
-P:databasenode.\"Source_ABC\".datasource="{\"ORA 10gR2\",  
user1, mypsw, true}"
```

以下の TM1 データ ソースの例のように、構造化パラメーターを示す引用符の前には円記号も必要です。

```
clemb -server -hostname 9.115.21.169 -port 28053 -username administrator  
-execute -stream C:\Share\TM1_Script.str -P:tm1import.pm_host="http://9.115.21.163:9510/  
pmhub/pm"  
-P:tm1import.tm1_connection="{\"SData\", \"\", \"admin\", \"apple\"}"  
-P:tm1import.selected_view="{\"SalesPriorCube\", \"salesmargin%\"}"
```

注: データベース名 (datasource プロパティ内) に 1 つ以上のスペース、ピリオド (「終止符」とも呼ばれる)、または下線が含まれる場合は、「円記号と二重引用符」形式を使用して、それを文字列として扱うことができます。例えば、"{\"db2v9.7.6_linux\"}" または "{\"TDATA 131\"}" です。さらに、次の例のように、常に datasource ストリング値を二重引用符と中括弧で囲みます。"{\"SQL Server\", spssuser, abcd1234, false}"。

サーバー接続の引数

-server フラグは、パブリック・サーバーに接続する必要があることを IBM スポス モデラーに通知し、IBM スポス モデラーにパブリック・サーバーへの接続方法を指示するためにフラグ -hostname、-use_ssl、-port、-username、-password、および -domain が使用されます。-server 引数が指定されていない場合、デフォルト・サーバーまたはローカル・サーバーが使用されます。

例

パブリック・サーバーに接続するには

```
modelerclient -server -hostname myserver -port 80 -username dminer  
-password 1234 -stream mystream.str -execute
```

サーバー・クラスターに接続するには

```
modelerclient -server -cluster "QA Machines" \  
-spsscr_hostname pes_host -spsscr_port 8080 \  
-spsscr_username asmith -spsscr_epassword xyz
```

サーバー・クラスターへの接続には IBM スポス コラボレーションおよびデプロイメント・サービスを介した Coordinator of Processes が必要であるため、-cluster 引数はリポジトリ接続オプション (spsscr_*) と組み合わせて使用が必要であることに注意してください。詳しくは、69 ページの『[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 接続の引数](#)』のトピックを参照してください。

表 36. サーバー接続の引数	
引数	動作説明
-server	IBM スポス モデラー をサーバー・モードで実行し、フラグ -hostname、-port、-username、-password、および -domain を使用してパブリック・サーバーに接続します。
-hostname <名>	サーバー・マシンのホスト名を指定します。サーバー・モードでしか利用できません。
-use_ssl	接続で使用する SSL (secure socket layer) を指定します。このフラグはオプションです。SSL 使用時のデフォルト設定は <i>not</i> です。
-port <番号>	指定したサーバーのポート番号。サーバー・モードでしか利用できません。
-cluster <名>	名前付きサーバーではなく、サーバー・クラスターへの接続を指定します。この引数は hostname、port、および use_ssl 引数の代替です。name はクラスター名、または 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・リポジトリ 内のクラスターを識別する一意の URI です。サーバー・クラスターは、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 を使用して Coordinator of Processes で管理されます。詳しくは、トピック 69 ページの『 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 接続の引数 』を参照してください。
-username <名>	サーバーにログオンするためのユーザー名。サーバー・モードでしか利用できません。
-password <password>	サーバーにログオンするためのパスワード。サーバー・モードでしか利用できません。 注: -password 引数を使用しない場合、パスワードの入力を要求するプロンプトが表示されます。
-epassword <encodedpasswordstring>	サーバーにログオンするための暗号化パスワード。サーバー・モードでしか利用できません。 注: 暗号化パスワードは、IBM スポス モデラー アプリケーションの「ツール」メニューから生成することができます。
-domain <名>	サーバーにログオンする際に使用するドメイン名。サーバー・モードでしか利用できません。
-P <名> = <value>	スタートアップ・パラメーターの設定に使用されます。ノードのプロパティ (スロット・パラメーター) の設定に使用することもできます。

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・リポジトリ 接続の引数

コマンド・ラインを経由して 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 でオブジェクトを保存したり取り出したりするには、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・リポジトリに有効な接続を指定する必要があります。例：

```
modelerclient -spsscr_hostname myhost -spsscr_port 8080
-spsscr_username myusername -spsscr_password mypassword
-stream "spsscr:///folder_1/scoring_stream.str" -execute
```

接続を設定するために使用できる引数の一覧を次の表に示します。

表 37. 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 接続の引数	
引数	動作説明
-spsscr_hostname <hostname or IP address>	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 がインストールされているサーバーのホスト名または IP アドレスです。
-spsscr_port <number>	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 が接続を承認したポート番号です（通常、8080 がデフォルト値）。
-spsscr_use_ssl	接続で使用する SSL (secure socket layer) を指定します。このフラグはオプションです。SSL 使用時のデフォルト設定は <i>not</i> です。
-spsscr_username <name>	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 にログオンするためのユーザー名。
-spsscr_password <password>	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 にログオンするためのパスワード。
-spsscr_epassword <encoded password>	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 にログオンするためのエンコードされたパスワード。
-spsscr_providername <name>	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 (Active Directory または LDAP) へのログオンに使用する認証プロバイダー。これは、ネイティブ (ローカル・リポジトリ) のプロバイダーを使用する場合は不要です。

IBMSPSS Analytic Server 接続の引数

コマンドラインを使用して IBMSPSS Analytic Server でオブジェクトを保存したり取り出したりするには、IBMSPSS Analytic Server への有効な接続を指定する必要があります。

注：分析サーバーのデフォルトの場所は、スポスモデラー・サーバーから取得されます。ユーザーは、「ツール」>「**Analytic Server 接続**」を使用して、独自の分析サーバー接続を定義することもできます。

接続を設定するために使用できる引数の一覧を次の表に示します。

表 38. IBMSPSS Analytic Server 接続の引数	
引数	動作説明
-analytic_server_username	IBMSPSS Analytic Server にログオンするためのユーザー名。
-analytic_server_password	IBMSPSS Analytic Server にログオンするためのパスワード。
-analytic_server_epassword	IBMSPSS Analytic Server にログオンするための暗号化パスワード。
-analytic_server_credential	IBMSPSS Analytic Server にログオンするために使用する資格情報。

複数の引数の組み合わせ

複数の引数を記述したコマンド・ファイルを作成し、起動時に @ 記号に続けてそのファイル名を指定することができます。こうすることによって、コマンド・ラインによる起動を短縮し、OS によるコマンド長の制限に関する問題を解決することができます。例えば、以下の起動コマンドは <commandFileName> が示すファイルに指定されている引数を使用します。

```
modelerclient @<commandFileName>
```


ファイル名やコマンド・ファイルへのパスにスペースがある場合は、以下のようにして引用符で囲みます。

```
modelerclient @ "C:\Program  
Files\IBM\SPSS\Modeler\nn\scripts\my_command_file.txt"
```

このコマンド・ファイルには、スタートアップ時に個別に指定していたすべての引数を記述することができます。以下に例を示します。

```
-stream report.str  
-Porder.full_filename=APR_orders.dat  
-Preport.filename=APR_report.txt  
-execute
```

コマンド・ファイルを記述して、コマンド・ファイル名を指定する場合の制限事項を次に示します。

- 1行につき1つの引数またはコマンドを記述する必要があります。
- コマンド・ファイル内に、@CommandFile 引数を組み込まないでください。

第7章 プロパティ・リファレンス

プロパティ・リファレンスの概要

ノード、ストリーム、プロジェクト、スーパーノードに対して、数多くのさまざまなプロパティを指定できます。名前、注釈、およびツールヒントなど、すべてのノードに共通のプロパティもありますが、その一方で、ノードのタイプに固有なプロパティもあります。キャッシングやスーパーノードの動作などの高レベルなストリーム操作を参照するプロパティもあります。プロパティは、標準のユーザー・インターフェースからアクセスでき(ノードのオプションを編集するダイアログ・ボックスを開く場合など)、また、多くの標準とは異なる方法でも使用できます。

- プロパティは、このセクションで説明されているように、スクリプトからアクセスできます。詳しくは、[73 ページの『プロパティのシンタックス』](#)を参照してください。
- ノードのプロパティは、スーパーノード・パラメーター中で使用することができます。
- ノード・プロパティは、IBM スポス モデラーの始動時にコマンド行オプションの一部として (-P フラグを使用して) 使用することもできます。

IBM スポス モデラー のスクリプトでは、ノードおよびストリームのプロパティは、よく **スロット・パラメーター** と呼ばれます。このガイドでは、スロット・パラメーターをノードまたはストリームのプロパティと記載しています。

プロパティのシンタックス

プロパティは、以下のシンタックスを使用して設定できます。

```
OBJECT.setPropertyValue(PROPERTY, VALUE)
```

または

```
OBJECT.setKeyedPropertyValue(PROPERTY, KEY, VALUE)
```

プロパティの値は、以下のシンタックスを使用して取得できます。

```
VARIABLE = OBJECT.getPropertyValue(PROPERTY)
```

または

```
VARIABLE = OBJECT.getKeyedPropertyValue(PROPERTY, KEY)
```

ここで、OBJECT はノードまたは出力、PROPERTY は式で参照しているノード プロパティの名前、KEY はキー プロパティのキー値です。例えば、以下のシンタックスを使用して、フィルター・ノードを検索し、すべてのフィールドを含むようにデフォルトを設定し、下流データから Age フィールドをフィルタリングします。

```
filternode = modeler.script.stream().findByType("filter", None)
filternode.setPropertyValue("default_include", True)
filternode.setKeyedPropertyValue("include", "Age", False)
```

IBM スポス モデラー で使用されるすべてのノードは、ストリームの `findByType(TYPE, LABEL)` 関数を使用して見つけることができます。少なくとも TYPE または LABEL のいずれかを指定する必要があります。

構造化プロパティー

スクリプト解析時の明確性を向上するために構造化プロパティーを使用するには、次の2種類の方法があります。

- データ型、フィルター、またはバランス・ノードなどの、複雑なノードのプロパティー名を構造化する。
- 複数のプロパティーを同時に指定する形式を提供する。

複雑なインターフェースの構造化

テーブルや他の複雑なインターフェースがあるノード、例えば、データ型、フィルター、およびバランス・ノードなどを対象とするスクリプトは、正しく解析されるために一定の構造を遵守する必要があります。これらの構造化プロパティーには、1つの識別子名と比べてより複雑な名前が必要です。これらのプロパティーでは、単一の識別子名よりも複雑な名前が必要であり、この名前はキーと呼ばれます。この情報を参照するため、フィルター・ノードではフィールドごとに1つの情報項目(各フィールドが真か偽か)が保存されます。この情報を参照するために、フィルター・ノードはフィールドごとに1つの情報項目を保管します(各フィールドが true か false か)。このプロパティーには、真 (True) または偽 (False) の値が設定されているか、または指定される可能性があります。mynode というフィルター・ノード (上流側) に、Age というフィールドがある場合を考えてみましょう。これをオフにするには、次のように、キー Age と値 False を指定してプロパティー include を設定します。

```
mynode.setKeyedPropertyValue("include", "Age", False)
```

複数のプロパティーの設定構造

多数のノードに対して、複数のノードおよびストリームのプロパティーを同時に割り当てることができます。これは、**multiset コマンド**または**セット ブロック**と呼ばれています。

場合によっては、構造化プロパティーがきわめて複雑なこともあります。以下に例を示します。

```
sortnode.setPropertyValue("keys", [["K", "Descending"], ["Age", "Ascending"], ["Na", "Descending"]])
```

構造化プロパティーのもう1つの利点は、ノードが安定していなくてもそのノード上に複数のプロパティーが設定できることです。デフォルトでは、multiset はブロック内のすべてのプロパティーを設定してから、個別のプロパティー設定に基づいてアクションを実行します。例えば固定長ファイル・ノードを定義するときに、フィールド・プロパティーを2ステップに分けて設定するとエラーが生じます。これは、両方の設定が有効になるまでノードが一貫しないためです。プロパティーを multiset として定義すれば、データ・モデルを更新する前に両方のプロパティーが設定でき、エラーが回避されます。

省略形

ノードのプロパティーのシンタックスでは、標準省略形が使用されています。省略形を覚えておけば、スクリプトの作成に役立ちます。

表 39. シンタックスで使用される標準省略形	
省略形	意味
abs	絶対値
len	Length
分	最小
最大	最大値
correl	相関
covar	共分散

表 39. シンタックスで使用する標準省略形 (続き)	
省略形	意味
num	数字または数値
pct	パーセントまたは割合
transp	透明度
xval	交差検証
var	分散または変数 (入力ノードで)

ノードおよびストリームのプロパティの例

ノードおよびストリームのプロパティは、IBM スポス モデラー のさまざまな場面で使用されます。一般的にこれらのプロパティは、複数のストリームや操作を自動化するために用いられる **スタンドアロン スクリプト**、または単一のストリーム内のプロセスの自動化に用いられる **ストリーム・スクリプト** など、スクリプトの一部として使われます。スーパーノード内で、ノードのプロパティを使用してノード・パラメーターを指定することもできます。もっとも基本的なレベルで、IBM スポス モデラー の起動時にコマンド・ライン・オプションとしてプロパティを指定することもできます。コマンド・ラインの起動時に、`-p` 引数を指定すれば、ストリーム・プロパティを使用してストリームの設定を変更することができます。

表 40. ノードおよびストリームのプロパティの例	
プロパティ	意味
<code>s.max_size</code>	ノード <code>s</code> のプロパティ <code>max_size</code> を表します。
<code>s:samplenode.max_size</code>	ノード <code>s</code> のプロパティ <code>max_size</code> を表します。 このノードは、サンプリング・ノードでなければなりません。
<code>:samplenode.max_size</code>	現在のストリーム中のサンプリング・ノードの、プロパティ <code>max_size</code> を表します (サンプリング・ノードは 1 つだけでなければなりません)。
<code>s:sample.max_size</code>	ノード <code>s</code> のプロパティ <code>max_size</code> を表します。 このノードは、サンプリング・ノードでなければなりません。
<code>t.direction.Age</code>	データ型ノード <code>t</code> の <code>Age</code> フィールドの役割を表します。
<code>:.max_size</code>	*** 無効 *** ノード名またはノードの種類を指定する必要があります。

`s:sample.max_size` の例は、ノードの種類を完全に記述する必要がないことを示しています。

`t.direction.Age` の例は、1 つのノードの属性が個別の値を持つ単純な個々のスロットよりも複雑な場合に、一部のスロット名を構造化できることを示しています。このようなスロットは、**構造化または複雑なプロパティ**と呼ばれます。

ノードのプロパティの概要

ノードの種類ごとに、独自の有効なプロパティのセットが用意されています。また、各プロパティにはデータ型があります。一般的なデータ型の数値、フラグ、または文字列の場合、プロパティの設定は強制的に正しいデータ型に設定されます。強制的に設定できない場合はエラーが発生します。それに対し、プロパティ参照が、`Discard`、`PairAndDiscard`、および `IncludeAsText` のような有効な値の範囲を指定していることもあります。この場合、範囲外の値が使われた場合にエラーになります。フラグ型プロパティは、`true` および `false` の値を使用して読み込まれるか、設定される必要があります (`Off`、

OFF、off、No、NO、no、n、N、f、F、false、False、FALSE、または0なども値の設定時に認識されますが、プロパティ値の読み込み時にエラーが発生する場合があります。その他の値はすべて真と見なされます。trueとfalseを使用すると、こうした混乱が避けられます)。このガイドにある参照テーブルでは、構造化プロパティはそのまま「**プロパティの説明**」欄に、使用形式とともに記載されています。

共通のノード・プロパティ

数多くのプロパティが、IBM スポス モデラー 中のすべてのノード (スーパーノードも含む) で共通に使われています。

表 41. 共通のノード・プロパティ		
プロパティ名	データ・タイプ	プロパティの説明
use_custom_name	フラグ	
name	string	ストリーム領域上のノード名を対象とする読み込み専用プロパティです (自動またはユーザー設定)。
custom_name	string	ノードのカスタム(ユーザー設定)名を指定します。
tooltip	string	
annotation	string	
keywords	string	オブジェクトに関連付けられているキーワードのリストを指定する構造化スロットです (例: ["Keyword1" "Keyword2"])
cache_enabled	フラグ	
node_type	source_supernode process_supernode terminal_supernode 指定されたすべてのノード名 スクリプト	ノードをタイプごとに参照するために使用される読み込み専用プロパティ。例えば、ノードを real_income のような名前だけで参照する代わりに、userinputnode または filternode のようなタイプで指定することもできます。

スーパーノード固有のプロパティは、他のノードと同様に、個別に説明します。詳しくは、トピック [451](#) ページの『[第 21 章 スーパーノードのプロパティ](#)』を参照してください。

第 8 章 Stream プロパティ

スクリプトにより、さまざまなストリームのプロパティを制御することができます。ストリームのプロパティを参照するには、以下のような、スクリプトを使用するための実行メソッドを設定する必要があります。

```
stream = modeler.script.stream()
stream.setPropertyValue("execute_method", "Script")
```

例

ノード プロパティを使用して、現在のストリーム内の各ノードが参照されます。次のストリーム・スクリプトに、その例を示します。

```
stream = modeler.script.stream()
annotation = stream.getPropertyValue("annotation")

annotation = annotation + "\n\nThis stream is called \"" + stream.getLabel() + "\" and
contains the following nodes:\n"

for node in stream.iterator():
    annotation = annotation + "\n" + node.getTypeName() + " node called \"" + node.getLabel()
    + "\""

stream.setPropertyValue("annotation", annotation)
```

この例では、ノード プロパティを使用して、ストリーム内のすべてのノードのリストを作成し、そのリストをストリームの注釈に書き込んでいます。この注釈は、次のようになります。

```
This stream is called "druglearn" and contains the following nodes:

type node called "Define Types"
derive node called "Na_to_K"
variablefile node called "DRUG1n"
neuralnetwork node called "Drug"
c50 node called "Drug"
filter node called "Discard Fields"
```

ストリームのプロパティを次の表に示します。

表 42. Stream プロパティ		
プロパティ名	データ・タイプ	プロパティの説明
execute_method	Normal	
	Script	

表 42. Stream プロパティ (続き)

プロパティ名	データ・タイプ	プロパティの説明
date_format	"DDMMYY" "MMDDYY" "YYMMDD" "YYYYMMDD" "YYYYDDD" DAY MONTH "DD-MM-YY" "DD-MM-YYYY" "MM-DD-YY" "MM-DD-YYYY" "DD-MON-YY" "DD-MON-YYYY" "YYYY-MM-DD" "DD.MM.YY" "DD.MM.YYYY" "MM.DD.YYYY" "DD.MON.YY" "DD.MON.YYYY" "DD/MM/YY" "DD/MM/YYYY" "MM/DD/YY" "MM/DD/YYYY" "DD/MON/YY" "DD/MON/YYYY" MON YYYY q Q YYYY ww WK YYYY	
date_baseline	number	
date_2digit_baseline	number	
time_format	"HHMMSS" "HHMM" "MMSS" "HH:MM:SS" "HH:MM" "MM:SS" "(H)H:(M)M:(S)S" "(H)H:(M)M" "(M)M:(S)S" "HH.MM.SS" "HH.MM" "MM.SS" "(H)H.(M)M.(S)S" "(H)H.(M)M" "(M)M.(S)S"	
time_rollover	フラグ	
import_datetime_as_string	フラグ	
decimal_places	number	
decimal_symbol	Default Period Comma	
angles_in_radians	フラグ	
use_max_set_size	フラグ	
max_set_size	number	

表 42. Stream プロパティ (続き)

プロパティ名	データ・タイプ	プロパティの説明
ruleset_evaluation	Voting FirstHit	
refresh_source_nodes	フラグ	ストリーム実行時に、入力ノードを自動的にリフレッシュするために使用します。
script	string	
annotation	string	
name	string	注: このプロパティは読み取り専用です。ストリーム名を変更する場合は、別名で保存する必要があります。
parameters		スタンドアロン スクリプト内からストリーム・パラメーターを更新する場合に、このプロパティを使用します。
nodes		詳細は以下を参照してください。
encoding	SystemDefault "UTF-8"	
stream_rewriting	Boolean	
stream_rewriting_maximise_sql	Boolean	
stream_rewriting_optimise_cl em_ execution	Boolean	
stream_rewriting_optimise_sy ntax_ execution	Boolean	
enable_parallelism	Boolean	
sql_generation	Boolean	
database_caching	Boolean	
sql_logging	Boolean	
sql_generation_logging	Boolean	
sql_log_native	Boolean	
sql_log_prettyprint	Boolean	
record_count_suppress_inpu t	Boolean	
record_count_feedback_int erval	integer	

表 42. Stream プロパティ (続き)

プロパティ名	データ・タイプ	プロパティの説明
use_stream_auto_create_node_settings	Boolean	true の場合はストリーム固有の設定が使用されます。それ以外の場合はユーザー設定が使用されます。
create_model_applier_for_new_models	Boolean	true の場合、モデル・ビルダーが新しいモデルを作成するときにアクティブな更新リンクがなければ、新しいモデル・アプ라이어が追加されます。 注: IBM スポス モデラー・バッチバージョン 15 を使用している場合は、スクリプト内で明示的にモデルアプライヤを追加する必要があります。
create_model_applier_update_links	createEnabled createDisabled doNotCreate	モデル・アプ라이어・ノードの自動追加時に作成するリンクの種類を定義します。
create_source_node_from_builders	Boolean	true の場合、ソース・ビルダーが新しいソース出力を作成するときにアクティブな更新リンクがなければ、新しい入力ノードが追加されます。
create_source_node_update_links	createEnabled createDisabled doNotCreate	入力ノードの自動追加時に作成するリンクの種類を定義します。
has_coordinate_system	Boolean	これを true に設定すると、ストリーム全体に座標系が適用されます。
coordinate_system	string	選択された投影座標系の名前。
deployment_area	ModelRefresh Scoring None	ストリームの展開方法を選択します。この値を None に設定すると、他の展開エントリーは使用されません。
scoring_terminal_node_id	string	ストリームのスコアリングブランチを選択します。ストリーム内のどのターミナル・ノードでもスコアリング・ブランチとして使用できます。
scoring_node_id	string	スコアリングブランチのナゲットを選択します。
model_build_node_id	string	ストリームのモデル作成ノードを選択します。

第9章 入力ノードのプロパティ

入力ノードの共通プロパティ

すべての入力ノードに共通するプロパティを次に一覧にします。その後に、特定のノードに関する情報が続きます。

例 1

```
varfilenode = modeler.script.stream().create("variablefile", "Var. File")
varfilenode.setPropertyValue("full_filename", "$CLEO_DEMOS/DRUG1n")
varfilenode.setKeyedPropertyValue("check", "Age", "None")
varfilenode.setKeyedPropertyValue("values", "Age", [1, 100])
varfilenode.setKeyedPropertyValue("type", "Age", "Range")
varfilenode.setKeyedPropertyValue("direction", "Age", "Input")
```

例 2

このスクリプトは、指定されたデータ ファイルに、複数行の文字列を表す Region というフィールドが含まれていることを前提とします。

```
from modeler.api import StorageType
from modeler.api import MeasureType

# Create a Variable File node that reads the data set containing
# the "Region" field
varfilenode = modeler.script.stream().create("variablefile", "My Geo Data")
varfilenode.setPropertyValue("full_filename", "C:/mydata/mygeodata.csv")
varfilenode.setPropertyValue("treat_square_brackets_as_lists", True)

# Override the storage type to be a list...
varfilenode.setKeyedPropertyValue("custom_storage_type", "Region",
StorageType.LIST)
# ...and specify the type if values in the list and the list depth
varfilenode.setKeyedPropertyValue("custom_list_storage_type", "Region",
StorageType.INTEGER)
varfilenode.setKeyedPropertyValue("custom_list_depth", "Region", 2)

# Now change the measurement to indentify the field as a geospatial value...
varfilenode.setKeyedPropertyValue("measure_type", "Region",
MeasureType.GEOSPATIAL)
# ...and finally specify the necessary information about the specific
# type of geospatial object
varfilenode.setKeyedPropertyValue("geo_type", "Region", "MultiLineString")
varfilenode.setKeyedPropertyValue("geo_coordinates", "Region", "2D")
varfilenode.setKeyedPropertyValue("has_coordinate_system", "Region", True)
varfilenode.setKeyedPropertyValue("coordinate_system", "Region",
"ETRS_1989_EPSG_Arctic_zone_5-47")
```

表 43. 入力ノードの共通プロパティー

プロパティー名	データ型	プロパティーの説明
direction	Input Target Both None Partition Split Frequency RecordID	フィールドの役割のキープロパティー。 使用形式： NODE.direction.FIELDNAME 注：値 In と Out は廃止されました。今後のリリースではサポートが中断される場合があります。
type	Range Flag Set Typeless Discrete Ordered Set Default	フィールドのデータ型。このプロパティーを <i>Default</i> に設定すると、 <i>values</i> プロパティーに関するすべての値は消去され、 <i>value_mode</i> を <i>Specify</i> に設定すると、それが <i>Read</i> にリセットされます。 <i>value_mode</i> が <i>Pass</i> または <i>Read</i> がすでに設定されている場合、 <i>type</i> の設定によって影響を受けることはありません。 使用形式： NODE.type.FIELDNAME
storage	Unknown String Integer Real Time Date Timestamp	フィールドのストレージ・タイプ用読み込み専用キー・プロパティー。 使用形式： NODE.storage.FIELDNAME

表 43. 入力ノードの共通プロパティ (続き)		
プロパティ名	データ型	プロパティの説明
check	None Nullify Coerce Discard Warn Abort	フィールド・タイプと範囲の検査用のキー・プロパティ。 使用形式： NODE.check.FIELDNAME
values	[値 値]	連続型 (範囲) フィールドの場合、最初の値が最小値で最後の値が最大値になります。名義型 (セット型) フィールドの場合、すべての値を指定します。フラグ型の場合、最初の値が <i>false</i> (偽) を、最後の値が <i>true</i> (真) を表します。このプロパティを設定すると、value_mode プロパティの値が自動的に <i>Specify</i> に設定されます。ストレージは、リストの最初の値に基づいて決まります。例えば、最初の値が文字列の場合、ストレージは String に設定されます。 使用形式： NODE.values.FIELDNAME
value_mode	Read Pass Read+ Current Specify	次のデータの受け渡し時にフィールドに値を設定する方法を決定します。 使用形式： NODE.value_mode.FIELDNAME このプロパティに <i>Specify</i> を直接には設定できないことに注意してください。特定の値を使用するには、values プロパティを設定します。
default_value_mode	Read Pass	すべてのフィールドに値を設定するためのデフォルトの方法を指定します。 使用形式： NODE.default_value_mode この設定による特定のフィールドの設定は、value_mode プロパティを使用するとオーバーライドされることがあります。

表 43. 入力ノードの共通プロパティ (続き)		
プロパティ名	データ型	プロパティの説明
extend_values	フラグ	value_mode が Read に設定された場合に適用されます。新しく読み込んだ値を、フィールドの既存の値に追加する場合は、Tを設定します。新しく読み込んだ値を優先して、既存の値を破棄する場合は、Fを設定します。 使用形式： NODE.extend_values.FIELDNAME
value_labels	string	値ラベルの指定に使用します。数値を先に指定します。
enable_missing	フラグ	Tを設定した場合、フィールドの欠損値の追跡が有効になります。 使用形式： NODE.enable_missing.FIELDNAME
missing_values	[value value ...]	欠損データを示すデータ値を指定します。 使用形式： NODE.missing_values.FIELDNAME
range_missing	フラグ	プロパティが T に設定されている場合、フィールドに欠損値 (空白) の範囲が定義されているかどうかを指定します。 使用形式： NODE.range_missing.FIELDNAME
missing_lower	string	range_missing が真 (true) の場合、欠損値範囲の下限値を指定します。 使用形式： NODE.missing_lower.FIELDNAME
missing_upper	string	range_missing が真 (true) の場合、欠損値範囲の上限値を指定します。 使用形式： NODE.missing_upper.FIELDNAME

表 43. 入力ノードの共通プロパティ (続き)		
プロパティ名	データ型	プロパティの説明
null_missing	フラグ	このプロパティが <i>T</i> に設定されていると、ヌル (ソフトウェアでは \$null\$ として表示される未定義値) は欠損値と見なされます。 使用形式: NODE.null_missing.FIELDNAME
whitespace_missing	フラグ	このプロパティが <i>T</i> に設定されていると、空白値 (スペース、タブ、および改行) だけを含む値は欠損値とみなされます。 使用形式: NODE.whitespace_missing.FIELDNAME
description	string	フィールドのラベルまたは説明の指定に使用します。
default_include	フラグ	デフォルトの処理としてフィールドを通過させるかフィルターをかけるかの指定をするキー・プロパティ。 NODE.default_include 例: set mynode:filternode.default_include = false
include	フラグ	各フィールドを適用するかフィルターをかけるかを決定するキー・プロパティ: NODE.include.FIELDNAME.
new_name	string	

表 43. 入力ノードの共通プロパティ (続き)

プロパティ名	データ型	プロパティの説明
measure_type	Range / MeasureType.RANGE Discrete / MeasureType.DISC ETE Flag / MeasureType.FLAG Set / MeasureType.SET OrderedSet / MeasureType.ORDER ED_SET Typeless / MeasureType.TYPEL ESS Collection / MeasureType.COLLE CTION Geospatial / MeasureType.GEOSP ATIAL	このキー付きプロパティは、フィールドに関連付けられた尺度を定義するために使用できるという点で、type と類似しています。異なるのは、Python スクリプトで、getter 関数が常に MeasureType 値を返す一方で、setter 関数に MeasureType 値のうちの 1 つを渡すこともできるという点です。
collection_measure	Range / MeasureType.RANGE Flag / MeasureType.FLAG Set / MeasureType.SET OrderedSet / MeasureType.ORDER ED_SET Typeless / MeasureType.TYPEL ESS	収集フィールド (深さが 0 のリスト) の場合、このキー付きプロパティは、基礎となる値に関連付けられた尺度タイプを定義します。

表 43. 入力ノードの共通プロパティ (続き)

プロパティ名	データ型	プロパティの説明
geo_type	Point MultiPoint LineString MultiLineString Polygon MultiPolygon	地理空間フィールドの場合、このキー付きプロパティにより、このフィールドが表す地理空間オブジェクトのタイプが定義されます。これは、値のリストの深さと整合している必要があります。
has_coordinate_system	Boolean	地理空間フィールドの場合、このプロパティにより、このフィールドに座標系があるかどうか定義されます。
coordinate_system	string	地理空間フィールドの場合、このキー付きプロパティにより、このフィールドの座標系が定義されます。
custom_storage_type	Unknown / MeasureType.UNKNOWN String / MeasureType.STRING Integer / MeasureType.INTEGER Real / MeasureType.REAL Time / MeasureType.TIME Date / MeasureType.DATE Timestamp / MeasureType.TIMESTAMP List / MeasureType.LIST	このキー付きプロパティは、フィールドのオーバーライドストレージを定義するために使用できるという点で、custom_storage と類似しています。異なるのは、Python スクリプトで、getter 関数が常に StorageType 値を返す一方で、setter 関数に StorageType 値のうちの 1 つを渡すこともできるという点です。

表 43. 入力ノードの共通プロパティ (続き)		
プロパティ名	データ型	プロパティの説明
custom_list_storage_type	String / MeasureType.STRING Integer / MeasureType.INTEGER Real / MeasureType.REAL Time / MeasureType.TIME Date / MeasureType.DATE Timestamp / MeasureType.TIMESTAMP	リスト フィールドの場合、このキー付きプロパティにより、基礎となる値のストレージ タイプが指定されます。
custom_list_depth	integer	リスト フィールドの場合、このキー付きプロパティにより、フィールドの深さが指定されます。
max_list_length	integer	地理空間 または集合 のいずれかの尺度を持つデータのみで使用できます。 リストの最大長を設定するには、リストに入れることができる要素の数を指定します。
max_string_length	integer	データ型不明 のデータでのみ使用可能で、SQL を生成してテーブルを作成するときに使用されます。データの最大文字列の値を入力します。これにより、テーブルに生成される列が、その文字列を含めるのに十分な大きさになります。

asimport プロパティ

分析サーバー 入力により、Hadoop 分散ファイル・システム (HDFS) でストリームを実行することができます。

例

```
node.setPropertyValue("use_default_as", False)
node.setPropertyValue("connection",
["false", "9.119.141.141", "9080", "analyticserver", "ibm", "admin", "admin", "false", "", "", "", "", ""])
```

表 44. asimport プロパティ		
asimport プロパティ	データ・タイプ	プロパティの説明
data_source	string	データ・ソースの名前。

表 44. asimport プロパティ (続き)

asimport プロパティ	データ・タイプ	プロパティの説明
use_default_as	Boolean	True に設定すると、サーバーの options.cfg ファイルで構成されているデフォルトの分析サーバー 接続が使用されます。False に設定した場合、このノードの接続が使用されます。
connection	["string", "string", "string", "string", "string", "string", "string", "string", "string", "string", "string"]	分析サーバー 接続の詳細を含むリストのプロパティ。形式は ["is_secure_connect", "server_url", "server_port", "context_root", "consumer", "user_name", "password", "use-kerberos-auth", "kerberos-krb5-config-file-path", "kerberos-jaas-config-file-path", "kerberos-krb5-service-principal-name", "enable-kerberos-debug"] です。ここで、is_secure_connect: はセキュア接続が使用されているかどうかを示し、true または false のいずれかです。use-kerberos-auth: は、Kerberos 認証が使用されるかどうかを示します。これは、true または false のいずれかです。enable-kerberos-debug: は、Kerberos 認証のデバッグ・モードが使用されているかどうかを示します。これは、true または false のいずれかです。

cognosimport ノードのプロパティ



IBM Cognos ソース・ノードは、Cognos Analytics データベースからデータをインポートします。

例

```
node = stream.create("cognosimport", "My node")
node.setPropertyValue("cognos_connection", ["http://mycogsrv1:9300/p2pd/
servlet/dispatch",
True, "", "", ""])
node.setPropertyValue("cognos_package_name", "/Public Folders/GOSALES")
node.setPropertyValue("cognos_items", "[[GreatOutdoors].[BRANCH].
[BRANCH_CODE]", "[GreatOutdoors]
.[BRANCH].[COUNTRY_CODE]]")
```

表 45. <i>cognosimport</i> ノードのプロパティ		
cognosimport ノードのプロパティ	データ・タイプ	プロパティの説明
mode	Data Report	Cognos データ (デフォルト) またはレポートをインポートするかどうかを指定します。

表 45. cognosimport ノードのプロパティ (続き)

cognosimport ノードのプロパティ	データ・タイプ	プロパティの説明
cognos_connection	["文字列", "フラグ", "文字列", "文字列", "文字列"]	Cognos サーバーの接続の詳細を含むリストのプロパティ。形式は以下のとおりです。["Cognos_server_URL", "login_mode", "namespace", "username", "password"] ここで、 Cognos_server_URL は、ソースが格納されている Cognos サーバーの URL です。 login_mode は、匿名ログインが使用されるかどうかを示します。true または false のいずれかです。true に設定されている場合、以下のフィールドを"." に設定する必要があります。 namespace はサーバーへのログオンに使用するセキュリティ認証プロバイダを示します。 username および password は Cognos サーバーにログオンする際に使用するユーザー名とパスワードです。 login_mode の代わりに、以下のモードも使用可能です。 - anonymousMode 以下に例を示します。["Cognos_server_url", "anonymousMode", "namespace", "username", "password"] - credentialMode 以下に例を示します。["Cognos_server_url", "credentialMode", "namespace", "username", "password"] - storedCredentialMode 以下に例を示します。["Cognos_server_url", "storedCredentialMode", "stored_credential_name"] ここで、stored_credential_name は、リポジトリ内での Cognos の資格情報の名前です。

表 45. *cognosimport* ノードのプロパティ (続き)

cognosimport ノードのプロパティ	データ・タイプ	プロパティの説明
cognos_package_name	string	データ・オブジェクトをインポートしている Cognos データ・ソース (通常はデータベース) のパスおよび名前。次に例を示します。 /Public Folders/GOSALES 注: スラッシュのみが有効です。
cognos_items	["フィールド", "フィールド", ..., "フィールド"]	インポートする 1 つまたは複数のデータ・オブジェクトの名前。フィールドの形式は、[namespace].[query_subject].[query_item] です。
cognos_filters	フィールド	データをインポートする前に適用するフィルターの名前。
cognos_data_parameters	リスト	データのプロンプト・パラメーターの値。名前と値のペアは大括弧で囲み、複数のペアはコンマで区切り、文字列全体は大括弧で囲みます。 フォーマット: [[<i>"param1"</i>, <i>"value"</i>], ..., [<i>"paramN"</i>, <i>"value"</i>]]
cognos_report_directory	フィールド	レポートをインポートするフォルダーまたはパッケージの Cognos パス。次に例を示します。 /Public Folders/GOSALES 注: スラッシュのみが有効です。
cognos_report_name	フィールド	インポートするレポートのレポートの位置内にあるパスと名前。
cognos_report_parameters	リスト	レポート・パラメーターの値。名前と値のペアは大括弧で囲み、複数のペアはコンマで区切り、文字列全体は大括弧で囲みます。 フォーマット: [[<i>"param1"</i>, <i>"value"</i>], ..., [<i>"paramN"</i>, <i>"value"</i>]]

databasenode プロパティ



データベース・ノードは、Microsoft SQL Server、Db2、Oracle など ODBC (開放型データベース接続) を使用する他のさまざまなパッケージからのデータのインポートに使用できます。

例

```
import modeler.api
stream = modeler.script.stream()
node = stream.create("database", "My node")
node.setPropertyValue("mode", "Table")
node.setPropertyValue("query", "SELECT * FROM drug1n")
node.setPropertyValue("datasource", "Drug1n_db")
node.setPropertyValue("username", "spss")
node.setPropertyValue("password", "spss")
node.setPropertyValue("tablename", ".Drug1n")
```

表 46. databasenode プロパティ

databasenode プロパティ	データ・タイプ	プロパティの説明
mode	Table Query	ダイアログ・ボックスのコントロールを使用してデータベースに接続するには、 <i>Table</i> を指定します。SQL を使用して選択されたデータベースにクエリーを行うには、 <i>Query</i> を指定します。
datasource	string	データベース名 (下記の注意を参照)。
username	string	データベース接続の詳細 (下記の注意を参照)。
password	string	
credential	string	IBM スポス コラボレーションおよびデプロイメント・サービスに保管されている資格情報の名前。このプロパティは、username プロパティや password プロパティの代わりに使用することができます。資格情報のユーザー名とパスワードは、データベースにアクセスするためのユーザー名とパスワードに一致している必要があります。
use_credential		True または False に設定します。
epassword	string	スクリプト内でパスワードをハードコード化する代わりに、エンコードされたパスワードを指定します。 詳しくは、トピック 54 ページの『暗号化パスワードの生成』 を参照してください。このプロパティは、実行時に読み取り専用になります。
tablename	string	アクセスするテーブルの名前。

表 46. databasenode プロパティ (続き)

databasenode プロパティ	データ・タイプ	プロパティの説明
strip_spaces	None Left Right Both	文字列の前後のスペースを破棄するためのオプションです。
use_quotes	AsNeeded Always Never	クエリーをデータベースに送信するときにテーブル名と列名を引用符で囲むかどうかを指定します (例えば、テーブル名と列名にスペースや句読点が含まれているような場合)。
query	string	送信するクエリーを表す SQL コードを指定します。

注: データベース名 (datasource プロパティ内) にスペースが含まれる場合、datasource、username、および password の個別のプロパティの代わりに、次の形式で単一のデータ ソース プロパティを使用することもできます。

表 47. databasenode プロパティ - datasource 固有

databasenode プロパティ	データ・タイプ	プロパティの説明
datasource	string	フォーマット: [database_name,username,password[,true false]] 暗号化パスワードと使用しないパラメーターです。true に設定すると、パスワードが使用前に復号されます。

データ・ソースを変更する場合、この形式を使用します。ただし、ユーザー名またはパスワードを変更する場合、username プロパティまたは password プロパティを使用できます。

datacollectionimportnode プロパティ



データ収集 データ・インポート・ノードで、市場調査製品によって使用される データ収集 Data Model に基づいた調査データをインポートします。このノードを使用するには、データ収集 Data Library がインストールされている必要があります。

例

```
node = stream.create("datacollectionimport", "My node")
node.setPropertyValue("metadata_name", "mrQvDsc")
node.setPropertyValue("metadata_file", "C:/Program Files/IBM/SPSS/
DataCollection/DDDL/Data/
Quanvert/Museum/museum.pkd")
node.setPropertyValue("casedata_name", "mrQvDsc")
node.setPropertyValue("casedata_source_type", "File")
node.setPropertyValue("casedata_file", "C:/Program Files/IBM/SPSS/
```



```
DataCollection/DDL/Data/
Quanvert/Museum/museum.pkd")
node.setPropertyValue("import_system_variables", "Common")
node.setPropertyValue("import_multi_response", "MultipleFlags")
```

表 48. datacollectionimportnode プロパティ

datacollectionimportnode プロパティ	データ・タイプ	プロパティの説明
metadata_name	string	MDSC の名前。特殊な値の DimensionsMDD は、標準的な データ収集 メタデータ・ドキュメントが使用される必要のあることを示します。ほかに、次の値を指定できます。 mrADODsc mrI2dDsc mrLogDsc mrQdiDrsDsc mrQvDsc mrSampleReportingMDSC mrSavDsc mrSCDsc mrScriptMDSC 特殊な値の none は、MDSC がないことを示します。
metadata_file	string	メタデータが格納されるファイルの名前。

表 48. *datacollectionimportnode* プロパティ (続き)

datacollectionimportnode プロパティ	データ・タイプ	プロパティの説明
<code>casedata_name</code>	<i>string</i>	CDSC の名前。可能な値は以下のとおりです。 <code>mrADODsc</code> <code>mrI2dDsc</code> <code>mrLogDsc</code> <code>mrPunchDSC</code> <code>mrQdiDrsDsc</code> <code>mrQvDsc</code> <code>mrRdbDsc2</code> <code>mrSavDsc</code> <code>mrScDSC</code> <code>mrXmlDsc</code> 特殊な値の <code>none</code> は、CDSC がないことを示します。
<code>casedata_source_type</code>	Unknown File Folder UDL DSN	CDSC のソース・タイプを示します。
<code>casedata_file</code>	<i>string</i>	<code>casedata_source_type</code> が <i>File</i> のときに、ケース・データが含まれるファイルを指定します。
<code>casedata_folder</code>	<i>string</i>	<code>casedata_source_type</code> が <i>Folder</i> のときに、ケース・データが含まれるフォルダーを指定します。
<code>casedata_udl_string</code>	<i>string</i>	<code>casedata_source_type</code> が <i>UDL</i> のときに、ケース・データが含まれるデータ・ソースのための OLD-DB 接続文字列を指定します。

表 48. datacollectionimportnode プロパティ (続き)

datacollectionimportnode プロパティ	データ・タイプ	プロパティの説明
casedata_dsn_string	string	casedata_source_type が DSN のときに、データ・ソースのための ODBC 接続文字列を指定します。
casedata_project	string	データ収集 データベースからケース・データを読み込むときに、プロジェクトの名前を入力できます。その他のケース・データのデータ型については、この設定を空白のままにしておく必要があります。
version_import_mode	All Latest Specify	各バージョンの取り扱い方法を定義します。
specific_version	string	version_import_mode が Specify のときに、インポートされるケース・データのバージョンを定義します。
use_language	string	特定言語のラベルが使用される必要があるかどうかを定義します。
language	string	use_language が真 (true) の場合、入力に使用する言語コードを定義します。言語コードは、ケース・データ内で利用できる中の 1 つにする必要があります。
use_context	string	特定のコンテキストが入力される必要があるかどうかを定義します。コンテキストは、応答に関連する説明を多様化させるために使用されます。
context	string	use_context が真 (true) の場合、入力するコンテキストを定義します。コンテキストは、ケース・データ内で利用できる中の 1 つにする必要があります。
use_label_type	string	特定のラベル タイプが入力される必要があるかどうかを定義します。
label_type	string	use_label_type が真 (true) の場合、入力するラベル・タイプを定義します。ラベル・タイプは、ケース・データ内で利用できる中の 1 つにする必要があります。
user_id	string	明示的なログインが必要なデータベースの場合、データ・ソースにアクセスするためのユーザー ID とパスワードを提供できます。
password	string	

表 48. datacollectionimportnode プロパティ (続き)

datacollectionimportnode プロパティ	データ・タイプ	プロパティの説明
import_system_variables	Common None All	インポートされるシステム変数を指定します。
import_codes_variables	フラグ	
import_sourcefile_variables	フラグ	
import_multi_response	MultipleFlags Single	

excelimportnode プロパティ



Excel インポート ノードは、Microsoft Excel から .xlsx ファイル形式でデータをインポートします。ODBC データ・ソースは不要です。

例

```
#To use a named range:
node = stream.create("excelimport", "My node")
node.setPropertyValue("excel_file_type", "Excel2007")
node.setPropertyValue("full_filename", "C:/drug.xlsx")
node.setPropertyValue("use_named_range", True)
node.setPropertyValue("named_range", "DRUG")
node.setPropertyValue("read_field_names", True)

#To use an explicit range:
node = stream.create("excelimport", "My node")
node.setPropertyValue("excel_file_type", "Excel2007")
node.setPropertyValue("full_filename", "C:/drug.xlsx")
node.setPropertyValue("worksheet_mode", "Name")
node.setPropertyValue("worksheet_name", "Drug")
node.setPropertyValue("explicit_range_start", "A1")
node.setPropertyValue("explicit_range_end", "F300")
```

表 49. excelimportnode プロパティ

excelimportnode プロパティ	データ・タイプ	プロパティの説明
excel_file_type	Excel2007	
full_filename	string	パスを含む、完全なファイル名。
use_named_range	Boolean	名前付けられた範囲を使用するかどうかを指定します。 真の場合、読み込む範囲を指定するのに named_range プロパティが使用され、その他のワークシートとデータ範囲の設定は無視されます。
named_range	string	

表 49. excelimportnode プロパティ (続き)

excelimportnode プロパティ	データ・タイプ	プロパティの説明
worksheet_mode	Index Name	ワークシートがインデックスで定義されているのか (Index)、または名前で定義されているのか (Name) を指定します。
worksheet_index	integer	読み込むべきワークシートのインデックス。最初のワークシートは 0、2 番目は 1、というようにインデックスが指します。
worksheet_name	string	読み込むべきワークシートの名前。
data_range_mode	FirstNonBlank ExplicitRange	範囲の決定方法を指定します。
blank_rows	StopReading ReturnBlankRows	data_range_mode が <i>FirstNonBlank</i> のときに、空白行の処理方法を指定します。
explicit_range_start	string	data_range_mode が <i>ExplicitRange</i> のときに、読み込む範囲の開始点を指定します。
explicit_range_end	string	
read_field_names	Boolean	指定された範囲の最初の行がフィールド (列) 名として使用されるかどうかを指定します。
scanLineCount	integer	列およびストレージ・タイプをスキャンする行の数を指定します。デフォルトは 200 です。

extensionimportnode プロパティ



拡張のインポート・ノードを使用すると、R スクリプトまたは Python for Spark スクリプトを実行して、データをインポートできます。

Python for Spark の例

```
##### Script example for Python for Spark
import modeler.api
stream = modeler.script.stream()
node = stream.create("extension_importer", "extension_importer")
node.setPropertyValue("syntax_type", "Python")

python_script = """
import spss.pyspark
from pyspark.sql.types import *

cxt = spss.pyspark.runtime.getContext()

_schema = StructType([StructField('id', LongType(), nullable=False), \
StructField('age', LongType(), nullable=True), \
StructField('Sex', StringType(), nullable=True), \
StructField('BP', StringType(), nullable=True), \
StructField('Cholesterol', StringType(), nullable=True), \
```

```

StructField('K', DoubleType(), nullable=True), \
StructField('Na', DoubleType(), nullable=True), \
StructField('Drug', StringType(), nullable=True)])

if cxt.isComputeDataModelOnly():
    cxt.setSparkOutputSchema(_schema)
else:
    df = cxt.getSparkInputData()
    if df is None:
        drugList=[(1,23,'F','HIGH','HIGH',0.792535,0.031258,'drugY'), \
(2,47,'M','LOW','HIGH',0.739309,0.056468,'drugC'),\
(3,47,'M','LOW','HIGH',0.697269,0.068944,'drugC'),\
(4,28,'F','NORMAL','HIGH',0.563682,0.072289,'drugX'),\
(5,61,'F','LOW','HIGH',0.559294,0.030998,'drugY'),\
(6,22,'F','NORMAL','HIGH',0.676901,0.078647,'drugX'),\
(7,49,'F','NORMAL','HIGH',0.789637,0.048518,'drugY'),\
(8,41,'M','LOW','HIGH',0.766635,0.069461,'drugC'),\
(9,60,'M','NORMAL','HIGH',0.777205,0.05123,'drugY'),\
(10,43,'M','LOW','NORMAL',0.526102,0.027164,'drugY')]
        sqlcxt = cxt.getSparkSQLContext()
        rdd = cxt.getSparkContext().parallelize(drugList)
        print 'pyspark read data count = '+str(rdd.count())
        df = sqlcxt.createDataFrame(rdd, _schema)

    cxt.setSparkOutputData(df)

node.setPropertyValue("python_syntax", python_script)

```

R の例

```

#### Script example for R
node.setPropertyValue("syntax_type", "R")

R_script = """# 'JSON Import' Node v1.0 for IBM SPSS Modeler
# 'RJSONIO' package created by Duncan Temple Lang - http://cran.r-project.org/web/packages/RJSONIO
# 'plyr' package created by Hadley Wickham http://cran.r-project.org/web/packages/plyr
# Node developer: Danil Savine - IBM Extreme Blue 2014
# Description: This node allows you to import into SPSS a table data from a JSON.
# Install function for packages
packages <- function(x){
  x <- as.character(match.call()[[2]])
  if (!require(x,character.only=TRUE)){
    install.packages(pkgs=x,repos="http://cran.r-project.org")
    require(x,character.only=TRUE)
  }
}
# packages
packages(RJSONIO)
packages(plyr)
#### This function is used to generate automatically the dataModel
getMetaData <- function (data) {
  if (dim(data)[1]<=0) {

    print("Warning : modelerData has no line, all fieldStorage fields set to strings")
    getStorage <- function(x){return("string")}

  } else {

    getStorage <- function(x) {
      res <- NULL
      #if x is a factor, typeof will return an integer so we treat the case on the side
      if(is.factor(x)) {
        res <- "string"
      } else {
        res <- switch(typeof(unlist(x)),
          integer = "integer",
          double = "real",
          character = "string",
          "string")
      }
      return (res)
    }
  }

  col = vector("list", dim(data)[2])
  for (i in 1:dim(data)[2]) {
    col[[i]] <- c(fieldName=names(data[i]),

```

```

        fieldLabel="",
        fieldStorage=getStorage(data[i]),
        fieldMeasure="",
        fieldFormat="",
        fieldRole="")
    }
    mdm<-do.call(cbind,col)
    mdm<-data.frame(mdm)
    return(mdm)
}
# From JSON to a list
txt <- readLines('C:/test.json')
formattedtxt <- paste(txt, collapse = '')
json.list <- fromJSON(formattedtxt)
# Apply path to json.list
if(strsplit(x='true', split='
',fixed=TRUE)[[1]][1]) {
  path.list <- unlist(strsplit(x='id_array', split=','))
  i = 1
  while(i<length(path.list)+1){
    if(is.null(getElement(json.list, path.list[i]))){
      json.list <- json.list[[1]]
    }else{
      json.list <- getElement(json.list, path.list[i])
      i <- i+1
    }
  }
}
# From list to dataframe via unlisted json
i <-1
filled <- data.frame()
while(i < length(json.list)+ 1){
  unlisted.json <- unlist(json.list[[i]])
  to.fill <- data.frame(t(as.data.frame(unlisted.json, row.names = names(unlisted.json))),
stringsAsFactors=FALSE)
  filled <- rbind.fill(filled,to.fill)
  i <- 1 + i
}
# Export to SPSS Modeler Data
modelerData <- filled
print(modelerData)
modelerDataModel <- getMetaData(modelerData)
print(modelerDataModel)
"""

node.setPropertyValue("r_syntax", R_script)

```

表 50. extensionimportnode プロパティ

extensionimportnode プロパティ	データ・タイプ	プロパティの説明
syntax_type	R <i>Python</i>	R または Python のどちらのスクリプトを実行するか指定します (R がデフォルトです)。
r_syntax	<i>string</i>	実行する R スクリプト・シンタックス。
python_syntax	<i>string</i>	実行する Python スクリプト・シンタックス。

fixedfilenode プロパティ



固定長ファイル・ノードで、固定長フィールド・テキスト・ファイルからデータをインポートします。ここで、ファイルのフィールドは区切られていませんが、同じ位置から始まって長さは固定されています。コンピュータによって生成されたデータや旧来のシステムのデータなどは、固定長フィールドの形式で保存されていることがよくあります。

例

```
node = stream.create("fixedfile", "My node")
node.setPropertyValue("full_filename", "$CLEO_DEMOS/DRUG1n")
node.setPropertyValue("record_len", 32)
node.setPropertyValue("skip_header", 1)
node.setPropertyValue("fields", [["Age", 1, 3], ["Sex", 5, 7], ["BP", 9,
10], ["Cholesterol",
12, 22], ["Na", 24, 25], ["K", 27, 27], ["Drug", 29, 32]])
node.setPropertyValue("decimal_symbol", "Period")
node.setPropertyValue("lines_to_scan", 30)
```

表 51. <i>fixedfilenode</i> プロパティ		
fixedfilenode プロパティ	データ・タイプ	プロパティの説明
record_len	<i>number</i>	各レコードの文字数を指定します。
line_oriented	フラグ	各レコードの末尾の改行文字をスキップします。
decimal_symbol	Default Comma Period	データ・ソースで使われている小数点記号。
skip_header	<i>number</i>	最初のレコードの先頭で無視する行数を指定します。列見出しを無視する場合などに役立ちます。
auto_recognize_datetime	フラグ	入力データの日付または時刻を自動的に特定するかどうかを指定します。
lines_to_scan	<i>number</i>	
fields	リスト	構造化プロパティ。
full_filename	<i>string</i>	読み込みファイルのディレクトリーを含む完全な名前。
strip_spaces	None Left Right Both	インポート時に文字列の前後のスペースを破棄します。
invalid_char_mode	Discard Replace	データ入力から不正な文字 (ヌル、0、または現在のエンコード中に存在していない文字) をデータ入力から削除するか (Discard)、指定された 1 文字の記号で不正な文字を置き換えます (Replace)。
invalid_char_replacement	<i>string</i>	
use_custom_values	フラグ	

表 51. *fixedfilenode* プロパティ (続き)

fixedfilenode プロパティ	データ・タイプ	プロパティの説明
custom_storage	Unknown String Integer Real Time Date Timestamp	

表 51. *fixedfilenode* プロパティ (続き)

fixedfilenode プロパティ	データ・タイプ	プロパティの説明
custom_date_format	"DDMMYY" "MMDDYY" "YYMMDD" "YYYYMMDD" "YYYYDDD" DAY MONTH "DD-MM-YY" "DD-MM-YYYY" "MM-DD-YY" "MM-DD-YYYY" "DD-MON-YY" "DD-MON-YYYY" "YYYY-MM-DD" "DD.MM.YY" "DD.MM.YYYY" "MM.DD.YY" "MM.DD.YYYY" "DD.MON.YY" "DD.MON.YYYY"	このプロパティは、カスタム (ユーザー設定) ストレージが指定される場合のみ適用されます。

表 51. *fixedfilenode* プロパティ (続き)

fixedfilenode プロパティ	データ・タイプ	プロパティの説明
	"DD/MM/YY" "DD/MM/YYYY" "MM/DD/YY" "MM/DD/YYYY" "DD/MON/YY" "DD/MON/YYYY" MON YYYY q Q YYYY ww WK YYYY	
custom_time_format	"HHMMSS" "HHMM" "MMSS" "HH:MM:SS" "HH:MM" "MM:SS" "(H)H:(M)M:(S)S" "(H)H:(M)M" "(M)M:(S)S" "HH.MM.SS" "HH.MM" "MM.SS" "(H)H.(M)M.(S)S" "(H)H.(M)M" "(M)M.(S)S"	このプロパティは、カスタム (ユーザー設定) ストレージが指定される場合のみ適用されます。

表 51. fixedfilenode プロパティ (続き)

fixedfilenode プロパティ	データ・タイプ	プロパティの説明
custom_decimal_symbol	フィールド	カスタム (ユーザー設定) ストレージが指定される場合のみ適用されます。
encoding	StreamDefault SystemDefault "UTF-8"	テキストのエンコード方法を指定します。

gsdata_import ノードのプロパティ



データマイニングセッションにマップまたは地理空間データを導入するには、地理空間入力ノードを使用します。

表 52. gsdata_import ノードのプロパティ

gsdata_import ノードのプロパティ	データ・タイプ	プロパティの説明
full_filename	string	ロードしたい .shp ファイルのパスを入力します。
map_service_URL	string	接続先のマップサービスの URL を入力します。
map_name	string	このプロパティには、マップサービスの最上位のフォルダー構造が格納されます (map_service_URL を使用する場合のみ)。

jsonimportnode のプロパティ



JSON 入力ノードは、JSON ファイルからデータをインポートします。

表 53. jsonimportnode のプロパティ

jsonimportnode のプロパティ	データ型	プロパティの説明
full_filename	string	パスを含む、完全なファイル名。
string_format	レコード 値	JSON スtring のフォーマットを指定します。デフォルトは records です。
auto_label		バージョン 18.2.1.1 で追加されました。

sasimportnode プロパティ



SAS インポート・ノードで、SAS データを IBM スポス モデラー ヘインポートします。

例

```
node = stream.create("sasimport", "My node")
node.setPropertyValue("format", "Windows")
node.setPropertyValue("full_filename", "C:/data/retail.sas7bdat")
node.setPropertyValue("member_name", "Test")
node.setPropertyValue("read_formats", False)
node.setPropertyValue("full_format_filename", "Test")
node.setPropertyValue("import_names", True)
```

表 54. sasimportnode プロパティ

sasimportnode プロパティ	データ・タイプ	プロパティの説明
format	Windows UNIX Transport SAS7 SAS8 SAS9	インポートするファイルの形式。
full_filename	string	パスも含めた、完全なファイル名。この名前を入力します。
member_name	string	指定した SAS トランスポート・ファイルからインポートするメンバーを指定します。
read_formats	フラグ	指定された形式ファイルから、データ形式 (変数ラベルなど) を読み込みます。
full_format_filename	string	
import_names	NamesAndLabels LabelsasNames	インポート時に変数名と変数ラベルをマッピングする方法を指定します。

simgennode プロパティ



シミュレーション生成ノードにより、シミュレーション対象のデータを容易に生成することができます。このとき、ユーザー指定の統計分布を使用して最初から生成するか、既存の履歴データに対してシミュレーションの当てはめノードを実行して得られた分布を使用して自動的に生成することができます。これは、モデル入力に不確実性が存在する場合に、予測モデルの結果を評価するときに役立ちます。

表 55. <i>simgen</i> node プロパティ		
simgen node プロパティ	データ・タイプ	プロパティの説明
fields	構造化プロパティ	例を参照
correlations	構造化プロパティ	例を参照
keep_min_max_setting	<i>Boolean</i>	
refit_correlations	<i>Boolean</i>	
max_cases	<i>integer</i>	最小値は 1000、最大値は 2,147,483,647 です。
create_iteration_field	<i>Boolean</i>	
iteration_field_name	<i>string</i>	
replicate_results	<i>Boolean</i>	
random_seed	<i>integer</i>	
parameter_xml	<i>string</i>	パラメーター XML を文字列として返します。

fields の例

これは、以下の構文を使用する構造化されたスロット パラメータです。

```
simgennode.setPropertyValue("fields", [
    [field1, storage, locked, [distribution1], min, max],
    [field2, storage, locked, [distribution2], min, max],
    [field3, storage, locked, [distribution3], min, max]
])
```

distribution は、分布名の宣言と、それに続く、属性の名前と値のペアを含むリストです。各分布は次のように定義されます。

```
[distributionname, [[par1], [par2], [par3]]]

simgennode = modeler.script.stream().createAt("simgen", u"Sim Gen", 726, 322)
simgennode.setPropertyValue("fields", [[["Age", "integer", False, ["Uniform", ["min", "1"],
["max", "2"]]]], "", ""]])
```

例えば、二項分布の単一フィールドを生成するノードを作成するために、以下のスクリプトを使用する場合があります。

```
simgen_node1 = modeler.script.stream().createAt("simgen", u"Sim Gen", 200, 200)
simgen_node1.setPropertyValue("fields", [[["Education", "Real", False, ["Binomial", ["n", 32],
["prob", 0.7]]], "", ""]])
```

二項分布では、**n** と **prob** の 2 つのパラメーターを使用します。二項分布では、最小値と最大値はサポートされず、空文字列として渡されます。

注: **distribution** を直接設定することはできません。これは、**fields** プロパティとともに使用します。

以下の例では、考えられるすべての分布タイプを示します。 **NegativeBinomialFailures** と **NegativeBinomialTrial** の両方でしきい値が **thresh** として入力されていることに注意してください。

```
stream = modeler.script.stream()
simgennode = stream.createAt("simgen", u"Sim Gen", 200, 200)

beta_dist = ["Field1", "Real", False, ["Beta", ["shape1", "1"], ["shape2", "2"]], "", ""]
binomial_dist = ["Field2", "Real", False, ["Binomial", ["n", "1"], ["prob", "1"]], "", ""]
categorical_dist = ["Field3", "String", False, ["Categorical", ["A", 0.3], ["B", 0.5], ["C", 0.2]], "", ""]
dice_dist = ["Field4", "Real", False, ["Dice", ["1", "0.5"], ["2", "0.5"]], "", ""]
```

```

exponential_dist = ["Field5", "Real", False, ["Exponential", [{"scale", "1"}]], "", ""]
fixed_dist = ["Field6", "Real", False, ["Fixed", [{"value", "1"}]], "", ""]
gamma_dist = ["Field7", "Real", False, ["Gamma", [{"scale", "1"}, {"shape", "1"}]], "", ""]
lognormal_dist = ["Field8", "Real", False, ["Lognormal", [{"a", "1"}, {"b", "1"}]], "", ""]
negbinomialfailures_dist = ["Field9", "Real", False, ["NegativeBinomialFailures", [{"prob", "0.5"}, {"thresh", "1"}]], "", ""]
negbinomialtrial_dist = ["Field10", "Real", False, ["NegativeBinomialTrials", [{"prob", "0.2"}, {"thresh", "1"}]], "", ""]
normal_dist = ["Field11", "Real", False, ["Normal", [{"mean", "1"}, {"stddev", "2"}]], "", ""]
poisson_dist = ["Field12", "Real", False, ["Poisson", [{"mean", "1"}]], "", ""]
range_dist = ["Field13", "Real", False, ["Range", [{"BEGIN", "1,3"}, {"END", "2,4"}], [{"PROB", "[0.5],[0.5]}"]], "", ""]
triangular_dist = ["Field14", "Real", False, ["Triangular", [{"min", "0"}, {"max", "1"}, {"mode", "1"}]], "", ""]
uniform_dist = ["Field15", "Real", False, ["Uniform", [{"min", "1"}, {"max", "2"}]], "", ""]
weibull_dist = ["Field16", "Real", False, ["Weibull", [{"a", "0"}, {"b", "1"}, {"c", "1"}]], "", ""]

```

```

simgennode.setPropertyValue("fields", [
  beta_dist, \
  binomial_dist, \
  categorical_dist, \
  dice_dist, \
  exponential_dist, \
  fixed_dist, \
  gamma_dist, \
  lognormal_dist, \
  negbinomialfailures_dist, \
  negbinomialtrial_dist, \
  normal_dist, \
  poisson_dist, \
  range_dist, \
  triangular_dist, \
  uniform_dist, \
  weibull_dist
])

```

correlations の例

これは、以下の構文を使用する構造化されたスロット パラメータです。

```

simgennode.setPropertyValue("correlations", [
  [field1, field2, correlation],
  [field1, field3, correlation],
  [field2, field3, correlation]
])

```

相関は、+1 から -1 までの任意の数字です。相関は必要な数だけ指定することができます。指定されていない相関は、すべて 0 に設定されます。不明なフィールドが存在する場合、相関値は相関行列 (または表) 上で設定する必要がある、赤いテキストで表示されます。不明なフィールドが存在する場合、ノードを実行することはできません。

statisticsimportnode プロパティ



IBM SPSS 統計 ファイル・ノードは、IBM SPSS 統計で使用される .sav ファイル・フォーマットからデータを読み取ります。また、IBM スポス モデラーに保存されているキャッシュ・ファイルも同じフォーマットを使用します。

このノードのプロパティについては、[423 ページの『statisticsimportnode プロパティ』](#)に記載されています。

tm1odataimport ノードのプロパティ



IBM Cognos TM1 入力ノードは、Cognos TM1 データベースからデータをインポートします。

表 56. tm1odataimport ノードのプロパティ

tm1odataimport ノードのプロパティ	データ・タイプ	プロパティの説明
credential_type	inputCredential または storedCredential	資格情報のタイプを示すために使用されます。

表 56. *tm1odataimport* ノードのプロパティ (続き)

tm1odataimport ノードのプロパティ	データ・タイプ	プロパティの説明
<code>input_credential</code>	リスト	<code>credential_type</code> が <code>inputCredential</code> のときは、ドメイン、ユーザー名、およびパスワードを指定します。
<code>stored_credential_name</code>	<i>string</i>	<code>credential_type</code> が <code>storedCredential</code> の場合、C & DS サーバー上の資格情報の名前を指定します。
<code>selected_view</code>	<code>["フィールド" "フィールド"]</code>	選択された TM1 キューブの詳細と、SPSS へのデータのインポートを行うキューブビューの名前を含むリストのプロパティ。以下に例を示します。 <code>TM1_import.setPropertyValue("selected_view", ['plan_BudgetPlan', 'Goal Input'])</code>
<code>is_private_view</code>	フラグ	<code>selected_view</code> が専用ビューであるかどうかを指定します。デフォルト値は <code>false</code> です。
<code>selected_columns</code>	<code>["フィールド"]</code>	選択した列を指定します。指定できる項目は 1 つのみです。以下に例を示します。 <code>setProperty("selected_columns", ["Measures"])</code>
<code>selected_rows</code>	<code>["フィールド" "フィールド"]</code>	選択した行を指定します。 例: <code>setProperty("selected_rows", ["Dimension_1_1", "Dimension_2_1", "Dimension_3_1", "Periods"])</code>
<code>connection_type</code>	<code>AdminServer</code> <code>TM1Server</code>	接続タイプを指定します。デフォルトは <code>AdminServer</code> です。
<code>admin_host</code>	<i>string</i>	REST API のホスト名の URL。 <code>connection_type</code> が <code>AdminServer</code> の場合は必須です。
<code>server_name</code>	<i>string</i>	<code>admin_host</code> から選択した TM1 サーバーの名前。 <code>connection_type</code> が <code>AdminServer</code> の場合は必須です。
<code>server_url</code>	<i>string</i>	TM1 サーバー REST API の URL。 <code>connection_type</code> が <code>TM1Server</code> の場合は必須です。

tm1import ノードのプロパティ (廃止)



IBM Cognos TM1 入力ノードは、Cognos TM1 データベースからデータをインポートします。

注: このノードは、Modeler 18.0 で廃止されました。それに置き換わるノードのスクリプト名は *tm1odataimport* です。

表 57. tm1import ノードのプロパティ

tm1import ノードのプロパティ	データ・タイプ	プロパティの説明
pm_host	string	注: バージョン 16.0 および 17.0 の場合のみ ホスト名。例: TM1_import.setPropertyValue("pm_host", 'http://9.191.86.82:9510/pmhub/pm')
tm1_connection	["フィールド", "フィールド", ..., "フィールド"]	注: バージョン 16.0 および 17.0 の場合のみ TM1 サーバーの接続の詳細を含むリストのプロパティ。形式は次のとおりです。 ["TM1_Server_Name", "tm1_username", "tm1_password"] 例: TM1_import.setPropertyValue("tm1_connection", ['Planning Sample', "admin", "apple"])
selected_view	["フィールド" "フィールド"]	選択された TM1 キューブの詳細と、SPSS へのデータのインポートを行うキューブビューの名前を含むリストのプロパティ。例: TM1_import.setPropertyValue("selected_view", ['plan_BudgetPlan', 'Goal Input'])
selected_column	["フィールド"]	選択した列を指定します。指定できる項目は 1 つのみです。 例: setPropertyValue("selected_columns", ["Measures"])
selected_rows	["フィールド" "フィールド"]	選択した行を指定します。 例: setPropertyValue("selected_rows", ["Dimension_1_1", "Dimension_2_1", "Dimension_3_1", "Periods"])

twcimport ノードのプロパティ



TWC 入力ノードは、IBM ビジネスの 1 つである The Weather Company から気象データをインポートします。これを使用して、ある場所の過去または予報の気象データを取得できます。これにより、使用可能な最も正確で高精度の気象データを利用して、意思決定を向上させるための気象主導のビジネス・ソリューションの開発に役立てることができます。

表 58. twcimport ノードのプロパティ

twcimport ノードのプロパティ	データ・タイプ	プロパティの説明
TWCDDataImport.latitude	real	緯度の値を形式 [-90.0~90.0] で指定します。
TWCDDataImport.longitude	real	経度の値を形式 [-180.0~180.0] で指定します。
TWCDDataImport.licenseKey	string	The Weather Company から入手したライセンス・キーを指定します。
TWCDDataImport.measurementUnit	English Metric Hybrid	測定単位を指定します。指定できる値は、English、Metric、Hybrid です。Metric がデフォルトです。
TWCDDataImport.dataType	Historical Forecast	入力する気象データのタイプを指定します。指定できる値は、Historical または Forecast です。Historical がデフォルトです。
TWCDDataImport.startDate	integer	TWCDDataImport.dataType に Historical を指定した場合は、開始日を yyyyMMdd の形式で指定します。
TWCDDataImport.endDate	integer	TWCDDataImport.dataType に Historical を指定した場合は、終了日を yyyyMMdd の形式で指定します。
TWCDDataImport.forecastHour	6 12 24 48	TWCDDataImport.dataType に Forecast を指定した場合は、時間に対して 6、12、24、または 48 を指定します。

userinputnode プロパティ



ユーザー入力ノードを利用すれば、最初から、あるいは既存のデータを変更して、合成データを簡単に作成できます。これは、モデル作成用の検定データセットを作成する場合などに役立ちます。

例

```
node = stream.create("userinput", "My node")
node.setPropertyValue("names", ["test1", "test2"])
node.setKeyedPropertyValue("data", "test1", "2, 4, 8")
node.setKeyedPropertyValue("custom_storage", "test1", "Integer")
node.setPropertyValue("data_mode", "Ordered")
```

表 59. *userinputnode* プロパティ

userinputnode プロパティ	データ・タイプ	プロパティの説明
data		
names		ノードにより生成されたフィールド名のリストを設定または返す構造化スロット。
custom_storage	Unknown String Integer Real Time Date Timestamp	フィールドのストレージを設定するか返す、キー・スロット。
data_mode	Combined Ordered	Combined が指定された場合、レコードは、セット値と最小/最大値のそれぞれ組み合わせについて生成されます。生成されたレコード数は、それぞれのフィールドの値の数値の積に等しくなります。Ordered が指定された場合、データ行を生成するために、各レコードの各列から 1 個の値が取られます。生成されるレコード数は、フィールドに関連付けられている最大の値に等しくなります。より小さいデータ値を持つフィールドは、ヌル値で埋められます。
values		注: このプロパティは <code>userinputnode.data</code> に置き換えられたため、使用しないでください。

variablefilenode プロパティ



可変長ファイル・ノードで、可変長フィールド・テキスト・ファイル、つまりフィールド数は一定でも各フィールド内の文字数が異なるレコードを含むファイルから、データを読み込みます。このノードは、固定長のヘッダー・テキストやある種の注釈があるファイルにも使用できます。

例

```
node = stream.create("variablefile", "My node")
node.setPropertyValue("full_filename", "$CLEO_DEMOS/DRUG1n")
node.setPropertyValue("read_field_names", True)
node.setPropertyValue("delimit_other", True)
node.setPropertyValue("other", ",")
node.setPropertyValue("quotes_1", "Discard")
node.setPropertyValue("decimal_symbol", "Comma")
node.setPropertyValue("invalid_char_mode", "Replace")
node.setPropertyValue("invalid_char_replacement", "|")
```

```
node.setKeyedPropertyValue("use_custom_values", "Age", True)
node.setKeyedPropertyValue("direction", "Age", "Input")
node.setKeyedPropertyValue("type", "Age", "Range")
node.setKeyedPropertyValue("values", "Age", [1, 100])
```

表 60. *variablefilenode* プロパティ

variablefilenode プロパティ	データ・タイプ	プロパティの説明
skip_header	<i>number</i>	最初のレコードの先頭で無視する文字数を指定します。
num_fields_auto	フラグ	各レコードのフィールドの数を自動的に決定します。レコードは、改行文字で終わる必要があります。
num_fields	<i>number</i>	各レコードのフィールドの数を手動で指定します。
delimit_space	フラグ	ファイルのフィールドを区切る文字を指定します。
delimit_tab	フラグ	
delimit_new_line	フラグ	
delimit_non_printing	フラグ	
delimit_comma	フラグ	この場合、コンマはストリーム内でフィールドの区切り文字と桁区切り記号の両方であるため、 <i>delimit_other</i> を <i>true</i> に設定し、 <i>other</i> プロパティを使用し、コンマを区切り記号として指定します。
delimit_other	フラグ	<i>other</i> プロパティを使用して、カスタム区切り記号をユーザーが指定できます。
other	<i>string</i>	<i>delimit_other</i> が <i>true</i> に設定されているときに使用される区切り記号を指定します。
decimal_symbol	Default Comma Period	データ・ソースで使われている小数点記号を指定します。
multi_blank	フラグ	複数の隣接するブランク区切り文字を 1 つの区切り文字として扱います。
read_field_names	フラグ	データ・ファイル中の最初の行を列のラベルとして取り扱います。
strip_spaces	None Left Right Both	インポート時に文字列の前後のスペースを破棄します。

表 60. *variablefilenode* プロパティ (続き)

variablefilenode プロパティ	データ・タイプ	プロパティの説明
<i>invalid_char_mode</i>	Discard Replace	データ入力から不正な文字 (ヌル、0、または現在のエンコード中に存在していない文字) をデータ入力から削除するか (Discard)、指定された 1 文字の記号で不正な文字を置き換えます (Replace)。
<i>invalid_char_replacement</i>	<i>string</i>	
<i>break_case_by_newline</i>	フラグ	行区切り文字が改行文字であることを指定します。
<i>lines_to_scan</i>	<i>number</i>	指定したデータ型をスキャンする行数を指定します。
<i>auto_recognize_datetime</i>	フラグ	入力データの日付または時刻を自動的に特定するかどうかを指定します。
<i>quotes_1</i>	Discard PairAndDiscard IncludeAsText	インポートでの単一引用符の処理方法を指定します。
<i>quotes_2</i>	Discard PairAndDiscard IncludeAsText	インポートでの二重引用符の処理方法を指定します。
<i>full_filename</i>	<i>string</i>	読み込みファイルのディレクトリーを含む完全な名前。
<i>use_custom_values</i>	フラグ	
<i>custom_storage</i>	Unknown String Integer Real Time Date Timestamp	

表 60. variablefilenode プロパティ (続き)

variablefilenode プロパティ	データ・タイプ	プロパティの説明
custom_date_format	"DDMMYY" "MMDDYY" "YYMMDD" "YYYYMMDD" "YYYYDDD" DAY MONTH "DD-MM-YY" "DD-MM-YYYY" "MM-DD-YY" "MM-DD-YYYY" "DD-MON-YY" "DD-MON-YYYY" "YYYY-MM-DD" "DD.MM.YY" "DD.MM.YYYY" "MM.DD.YY" "MM.DD.YYYY" "DD.MON.YY" "DD.MON.YYYY"	カスタム (ユーザー設定) ストレージが指定される場合のみ適用されます。

表 60. variablefilenode プロパティ (続き)		
variablefilenode プロパティ	データ・タイプ	プロパティの説明
	"DD/MM/YY" "DD/MM/YYYY" "MM/DD/YY" "MM/DD/YYYY" "DD/MON/YY" "DD/MON/YYYY" MON YYYY q Q YYYY ww WK YYYY	

表 60. variablefilenode プロパティ (続き)

variablefilenode プロパティ	データ・タイプ	プロパティの説明
custom_time_format	"HHMMSS" "HHMM" "MMSS" "HH:MM:SS" "HH:MM" "MM:SS" "(H)H:(M)M:(S)S" "(H)H:(M)M" "(M)M:(S)S" "HH.MM.SS" "HH.MM" "MM.SS" "(H)H.(M)M.(S)S" "(H)H.(M)M" "(M)M.(S)S"	カスタム (ユーザー設定) ストレージが指定される場合のみ適用されます。
custom_decimal_symbol	フィールド	カスタム (ユーザー設定) ストレージが指定される場合のみ適用されます。
encoding	StreamDefault SystemDefault "UTF-8"	テキストのエンコード方法を指定します。

xmlimportnode プロパティ



XML 入力ノードを使用して、XML 形式のデータをストリームにインポートできます。ディレクトリーの 1 つのファイルまたはすべてのファイルをインポートできます。オプションで、XML 構造を読み込むスキーマ ファイルを指定できます。

例

```
node = stream.create("xmlimport", "My node")
node.setPropertyValue("full_filename", "c:/import/ebooks.xml")
node.setPropertyValue("records", "/author/name")
```

表 61. *xmlimportnode* プロパティ

xmlimportnode プロパティ	データ・タイプ	プロパティの説明
read	single directory	単独のデータ・ファイルを読み込む (デフォルト) か、ディレクトリー内のすべての XML ファイルを読み込みます。
recurse	フラグ	指定したディレクトリーのすべてのサブディレクトリーから XML ファイルを追加で読み込むかどうかを指定します。
full_filename	string	(必須) インポートする XML ファイルの絶対パスとファイル名 (read = single の場合)。
directory_name	string	(必須) XML ファイルをインポートするディレクトリーの完全パスおよび名前 (read = directory の場合)。
full_schema_filename	string	XML 構造を読み込む XSD ファイルまたは DTD ファイルの完全パスおよびファイル名。このパラメーターを使用すると、構造を XML 入力ファイルから読み込みます。
records	string	レコードの境界を定義する XPath 式 (例: /author/name)。入力ファイルにこの要素が出現するごとに、新しいレコードが作成されます。
mode	read specify	すべてのデータを読み込む (デフォルト) か、読み込む項目を指定します。
fields		インポートする項目 (要素と属性) のリスト。リスト内の各アイテムは XPath 式です。

第 10 章 レコード設定ノードのプロパティ

appendnode プロパティ



レコード追加ノードで、レコードのセットを連結します。レコード追加ノードは、構造が似ていながらデータが異なるデータ・セットを組み合せる場合に役立ちます。

例

```
node = stream.create("append", "My node")
node.setPropertyValue("match_by", "Name")
node.setPropertyValue("match_case", True)
node.setPropertyValue("include_fields_from", "All")
node.setPropertyValue("create_tag_field", True)
node.setPropertyValue("tag_field_name", "Append_Flag")
```

表 62. appendnode プロパティ		
appendnode プロパティ	データ・タイプ	プロパティの説明
match_by	Position Name	メイン・データ・ソース中のフィールドの位置 (Position)、または入力データセット中のフィールド名 (Name) を基準にして、データセットを追加できます。
match_case	フラグ	フィールド名を比較するときに大文字と小文字の区別を有効にします。
include_fields_from	Main All	
create_tag_field	フラグ	
tag_field_name	string	

aggregatenode プロパティ



レコード集計ノードで、一連の入力レコードを要約集計された出力レコードに置き換えます。

例

```
node = stream.create("aggregate", "My node")
# dbnode is a configured database import node
stream.link(dbnode, node)
node.setPropertyValue("contiguous", True)
node.setPropertyValue("keys", ["Drug"])
node.setKeyedPropertyValue("aggregates", "Age", ["Sum", "Mean"])
node.setPropertyValue("inc_record_count", True)
node.setPropertyValue("count_field", "index")
node.setPropertyValue("extension", "Aggregated_")
node.setPropertyValue("add_as", "Prefix")
```

表 63. aggregatenode プロパティ

aggregatenode プロパティ	データ・タイプ	プロパティの説明
keys	リスト	集計にキーとして使用できるフィールドが一覧表示されます。例えば、キー・フィールドが Sex と Region の場合、一意な M と F の、および地域 N と S のそれぞれの組み合わせに対して集計レコードが作成されます (4 つの一意な組み合わせ)。
contiguous	フラグ	同じキー値を持つすべてのレコードが入力にグループ化されている場合 (例えば、入力にキー・フィールドにソートされる場合)、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、パフォーマンスが向上します。
aggregates		集計する数値フィールド、および選択されている集計モードを表示する構造化プロパティ。
aggregate_exprs		派生フィールドの名前を、そのフィールドを計算するために使用される集計式と共にキー化するキー プロパティ。以下に例を示します。 <pre>aggregatenode.setKeyedPropertyValue("aggregate_exprs", "Na_MAX", "MAX('Na')")</pre>
extension	string	重複集計フィールドに対応させる接頭辞または接尾辞を指定します (下の例を参照)。
add_as	Suffix Prefix	
inc_record_count	フラグ	各集計レコードを作成するために集計された入力レコード数を指定する追加フィールドを作成します。
count_field	string	レコード度数フィールドの名前を指定します。
allow_approximation	Boolean	分析サーバーでの集計の実行時に順序統計の近似を許可します。
bin_count	integer	近似で使用するビン数を指定します。

balancenode プロパティ



バランス・ノードで、データ・セットが指定した条件に合うように、データ・セットの不均衡を修正します。バランス式で、指定した比率によって条件が真 (true) の場合に、レコードの比率を調整します。

例

```
node = stream.create("balance", "My node")
node.setPropertyValue("training_data_only", True)
node.setPropertyValue("directives", [[1.3, "Age > 60"], [1.5, "Na > 0.5"]])
```

表 64. *balancenode* プロパティ

balancenode プロパティ	データ・タイプ	プロパティの説明
directives		指定された数値に基づいてフィールド値の割合を均衡にするための構造化プロパティ (次の例を参照してください)。
training_data_only	フラグ	学習データのみがバランス化されるよう指定します。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。

このノードのプロパティは次の形式を使用します。

[[数値、ストリング]\[数値、ストリング]\...[数値、ストリング]]。

注: ストリング (二重引用符を使用) が式に組み込まれている場合は、ストリングの前にエスケープ文字 " \ " を付ける必要があります。 " \ " 文字は、行継続文字でもあり、これを使用して、引数を明確に位置合わせすることができます。

cplexoptnode プロパティ



CPLEX の最適化ノードにより、OPL (Optimization Programming Language) モデル・ファイルを介した複雑な数学 (CPLEX) ベースの最適化の機能が提供されます。この機能は、IBM Analytical Decision Management 製品で使用できましたが、もうサポートされていません。ただし、IBM Analytical Decision Management を必要とせず、スポンサー モデラーで CPLEX ノードを使用することもできます。

表 65. *cplexoptnode* プロパティ

cplexoptnode プロパティ	データ型	プロパティの説明
opl_model_text	string	CPLEX の最適化ノードが実行し、最適化の結果を生成する OPL (Optimization Programming Language) スクリプト プログラム。
opl_tuple_set_name	string	入力データに対応する、OPL モデルのタプルセット名。これは必須ではなく、スクリプトによって設定されることは通常ありません。選択されたデータソースのフィールドマッピングを編集するためにのみ使用してください。
data_input_map	構造化プロパティのリスト	データソースの入力フィールドマッピング。これは必須ではなく、スクリプトによって設定されることは通常ありません。選択されたデータソースのフィールドマッピングを編集するためにのみ使用してください。

表 65. cplexoptnode プロパティ (続き)

cplexoptnode プロパティ	データ型	プロパティの説明
md_data_input_map	構造化プロパティのリスト	OPL 内で定義された各タプル間のフィールド マッピング、およびそれぞれに対応するフィールド データ ソース (入力データ)。ユーザーは、データ ソースごとにそれぞれを個別に編集できます。このスクリプトを使用すると、プロパティを直接設定して、すべてのマッピングを一度に設定できます。この設定は、ユーザー インターフェースに表示されません。 リスト内の各エンティティは構造化データです。 データ ソース タグ: データ ソースのタグ。これは、データ ソース ドロップダウンに表示されます。例えば、0_Products_Type の場合、タグは 0 です。 データ ソース インデックス: データ ソースの物理シーケンス (インデックス)。これは、接続順序に応じて決定されます。 ソース・ノード。 データ ソースの入力ノード (注釈)。これは、データ ソース ドロップダウンに表示されます。例えば、0_Products_Type の場合、入力ノードは Products です。 接続済みノード: 現在の CPLEX の最適化ノードに接続された事前ノード (注釈)。これは、データ ソース ドロップダウンに表示されます。例えば、0_Products_Type の場合、接続済みノードは Type です。 タプル セット名: データ・ソースのタプル セット名。これは、OPL 内の定義と一致する必要があります。 タプル フィールド名: データ ソースのタプル セット フィールド名。これは、OPL タプル セット定義内の定義と一致する必要があります。 ストレージ・タイプ: フィールド ストレージ タイプ。指定できる値は、int、float、または string です。

表 65. cplexoptnode プロパティ (続き)

cplexoptnode プロパティ	データ型	プロパティの説明
		データ フィールド名: データ・ソースのフィールド名。 例: <pre>[[0,0,'Product','Type','Products','prod_id_tup','int','prod_id'], [0,0,'Product','Type','Products','prod_name_tup','string','prod_name'], [1,1,'Components','Type','Components','comp_id_tup','int','comp_id'], [1,1,'Components','Type','Components','comp_name_tup','string','comp_name']]</pre>
opl_data_text	string	OPL に使用する変数、つまりデータの定義。
output_value_mode	string	指定できる値は、raw または dvar です。dvar を指定した場合、ユーザーは「出力」タブで OPL の目的関数の変数名を出力用に指定する必要があります。raw を指定した場合、名前に関係なく、目的関数が直接出力されます。
decision_variable_name	string	目的関数変数名は、OPL 内で定義されます。これは、output_value_mode プロパティが dvar に設定されている場合にのみ有効になります。
objective_function_value_fieldname	string	出力で使用する目的関数値のフィールド名。デフォルトは _OBJECTIVE です。
output_tuple_set_names	string	入力データからの事前定義されたタプルの名前。これは、決定変数のインデックスとしての役割を果たし、「変数の出力」で出力されることが予想されます。「出力タプル」は、OPL 内の決定変数定義と整合性を持つ必要があります。複数のインデックスがある場合、タプル名は、コンマ(,)で結合される必要があります。 単一タプルの例は Products で、対応する OPL 定義は dvar float+ Production[Products]; です。 複数のタプルの例は Products,Components で、対応する OPL 定義は dvar float+ Production[Products][Components]; です。

表 65. cplexoptnode プロパティ (続き)

cplexoptnode プロパティ	データ型	プロパティの説明
decision_output_map	構造化プロパティのリスト	OPL 内で定義された出力される変数と出力フィールドとの間のフィールド マッピング。リスト内の各エンティティは構造化データです。 変数名。 OPL 内の出力する変数名。 ストレージ・タイプ: 指定できる値は、int、float、または string です。 出力フィールド名: 結果 (出力またはエクスポート) 内の予期されるフィールド名。 例: <pre>[['Production', 'int', 'res'], ['Remark', 'string', 'res_1'] ['Cost', 'float', 'res_2']]</pre>

derive_stbnode プロパティ



スペース・タイム・ボックス・ノードは、緯度、経度、およびタイム・スタンプの各フィールドから、スペース-時間-ボックスを派生させます。頻度の高いスペース・タイム・ボックスをハングアウトとして識別することもできます。

例

```
node = modeler.script.stream().createAt("derive_stb", "My node", 96, 96)

# Individual Records mode
node.setPropertyValue("mode", "IndividualRecords")
node.setPropertyValue("latitude_field", "Latitude")
node.setPropertyValue("longitude_field", "Longitude")
node.setPropertyValue("timestamp_field", "OccurredAt")
node.setPropertyValue("densities", ["STB_GH7_1HOUR", "STB_GH7_30MINS"])
node.setPropertyValue("add_extension_as", "Prefix")
node.setPropertyValue("name_extension", "stb_")

# Hangouts mode
node.setPropertyValue("mode", "Hangouts")
node.setPropertyValue("hangout_density", "STB_GH7_30MINS")
node.setPropertyValue("id_field", "Event")
node.setPropertyValue("qualifying_duration", "30MINUTES")
node.setPropertyValue("min_events", 4)
node.setPropertyValue("qualifying_pct", 65)
```

表 66. スペース・タイム・ボックス・ノードのプロパティ

derive_stbnode プロパティ	データ・タイプ	プロパティの説明
mode	IndividualRecords Hangouts	
latitude_field	フィールド	
longitude_field	フィールド	

表 66. スペース・タイム・ボックス・ノードのプロパティ (続き)

derive_stbnode プロパティ	データ・タイプ	プロパティの説明
timestamp_field	フィールド	
hangout_density	密度	単一密度。有効な密度値については、「densities」を参照してください。
densities	[密度、密度、...、密度]	各 density は、STB_GH8_1DAY などの文字列です。 注: どの density が有効であるかについては、制約があります。geohash の場合、GH1 から GH15 の値を使用できます。この部分では、以下の値を使用できます <pre> EVER 1YEAR 1MONTH 1DAY 12HOURS 8HOURS 6HOURS 4HOURS 3HOURS 2HOURS 1HOUR 30MINS 15MINS 10MINS 5MINS 2MINS 1MIN 30SECS 15SECS 10SECS 5SECS 2SECS 1SEC </pre>
id_field	フィールド	
qualifying_duration	<pre> 1DAY 12HOURS 8HOURS 6HOURS 4HOURS 3HOURS 2Hours 1HOUR 30MIN 15MIN 10MIN 5MIN 2MIN 1MIN 30SECS 15SECS 10SECS 5SECS 2SECS 1SECS </pre>	文字列でなければなりません。
min_events	integer	最小の有効な整数値は 2 です。
qualifying_pct	integer	1 から 100 の範囲でなければなりません。
add_extension_as	<pre> Prefix Suffix </pre>	

表 66. スペース・タイム・ボックス・ノードのプロパティ (続き)

derive_stbnode プロパティ	データ・タイプ	プロパティの説明
name_extension	string	

distinctnode プロパティ



重複レコード・ノードで、重複レコードを削除します。その場合、最初の実重複するレコードをデータ・ストリームに渡すか、または、最初のレコードを破棄して、その後の重複レコードをデータ・ストリームに渡します。

例

```
node = stream.create("distinct", "My node")
node.setPropertyValue("mode", "Include")
node.setPropertyValue("fields", ["Age" "Sex"])
node.setPropertyValue("keys_pre_sorted", True)
```

表 67. distinctnode プロパティ

distinctnode プロパティ	データ・タイプ	プロパティの説明
mode	Include Discard	データ・ストリームに最初の実重複レコードを含めるか、最初の実重複レコードを破棄して、代わりにすべての重複レコードをデータ・ストリームに渡すことができます。
grouping_fields	リスト	レコードが同一であるかどうかを判断するために使われるフィールドを表示します。 注: このプロパティは、IBM スポス モデラー 16 以降では廃止されています。
composite_value	構造化スロット	下の例を参照してください。
composite_values	構造化スロット	下の例を参照してください。
inc_record_count	フラグ	各集計レコードを作成するために集計された入力レコード数を指定する追加フィールドを作成します。
count_field	string	レコード度数フィールドの名前を指定します。
sort_keys	構造化スロット。	注: このプロパティは、IBM スポス モデラー 16 以降では廃止されています。
default_ascending	フラグ	
low_distinct_key_count	フラグ	キー・フィールドに少ないレコードまたは少ない一意の値を持つよう指定します。
keys_pre_sorted	フラグ	同じキー値を持つすべてのレコードが入力で一緒にグループ化されるよう指定します。
disable_sql_generation	フラグ	

composite_value プロパティの例

composite_value プロパティは、以下の一般形式になっています。

```
node.setKeyedPropertyValue("composite_value", FIELD, FILLOPTION)
```

FILLOPTION の形式は[FillType, Option1, Option2, ...]です。

例:

```
node.setKeyedPropertyValue("composite_value", "Age", ["First"])
node.setKeyedPropertyValue("composite_value", "Age", ["last"])
node.setKeyedPropertyValue("composite_value", "Age", ["Total"])
node.setKeyedPropertyValue("composite_value", "Age", ["Average"])
node.setKeyedPropertyValue("composite_value", "Age", ["Min"])
node.setKeyedPropertyValue("composite_value", "Age", ["Max"])
node.setKeyedPropertyValue("composite_value", "Date", ["Earliest"])
node.setKeyedPropertyValue("composite_value", "Date", ["Latest"])
node.setKeyedPropertyValue("composite_value", "Code", ["FirstAlpha"])
node.setKeyedPropertyValue("composite_value", "Code", ["LastAlpha"])
```

カスタム オプションでは、複数の引数が必要であり、それらはリストとして追加されます。例えば、以下のようになります。

```
node.setKeyedPropertyValue("composite_value", "Name", ["MostFrequent", "FirstRecord"])
node.setKeyedPropertyValue("composite_value", "Date", ["LeastFrequent", "LastRecord"])
node.setKeyedPropertyValue("composite_value", "Pending", ["IncludesValue", "T", "F"])
node.setKeyedPropertyValue("composite_value", "Marital", ["FirstMatch", "Married", "Divorced", "Separated"])
node.setKeyedPropertyValue("composite_value", "Code", ["Concatenate"])
node.setKeyedPropertyValue("composite_value", "Code", ["Concatenate", "Space"])
node.setKeyedPropertyValue("composite_value", "Code", ["Concatenate", "Comma"])
node.setKeyedPropertyValue("composite_value", "Code", ["Concatenate", "UnderScore"])
```

composite_values プロパティの例

composite_values プロパティは、以下の一般形式になっています。

```
node.setPropertyValue("composite_values", [
    [FIELD1, [FILLOPTION1]],
    [FIELD2, [FILLOPTION2]],
    :
])
```

例:

```
node.setPropertyValue("composite_values", [
    ["Age", ["First"]],
    ["Name", ["MostFrequent", "First"]],
    ["Pending", ["IncludesValue", "T"]],
    ["Marital", ["FirstMatch", "Married", "Divorced", "Separated"]],
    ["Code", ["Concatenate", "Comma"]]
])
```

extensionprocessnode プロパティ



拡張の変換ノードを使用すると、R スクリプトまたは Python for Spark スクリプトを使用して、ストリームからデータを取得し、取得したデータに変換を適用できます。

Python for Spark の例

```
##### script example for Python for Spark
import modeler.api
stream = modeler.script.stream()
node = stream.create("extension_process", "extension_process")
node.setPropertyValue("syntax_type", "Python")

process_script = """
import spss.pyspark.runtime
from pyspark.sql.types import *

cxt = spss.pyspark.runtime.getContext()

if cxt.isComputeDataModelOnly():
    _schema = StructType([StructField("Age", LongType(), nullable=True), \
                             StructField("Sex", StringType(), nullable=True), \
                             StructField("BP", StringType(), nullable=True), \
                             StructField("Na", DoubleType(), nullable=True), \
                             StructField("K", DoubleType(), nullable=True), \
                             StructField("Drug", StringType(), nullable=True)])
    cxt.setSparkOutputSchema(_schema)
else:
    df = cxt.getSparkInputData()
    print df.dtypes[:]
    _newDF = df.select("Age", "Sex", "BP", "Na", "K", "Drug")
    print _newDF.dtypes[:]
    cxt.setSparkOutputData(_newDF)
"""

node.setPropertyValue("python_syntax", process_script)
```

R の例

```
##### script example for R
node.setPropertyValue("syntax_type", "R")
node.setPropertyValue("r_syntax", """"day<-as.Date(modelerData$dob, format="%Y-%m-%d")
next_day<-day + 1
modelerData<-cbind(modelerData,next_day)
var1<-c(fieldName="Next day",fieldLabel="",fieldStorage="date",fieldMeasure="",fieldFormat="",
fieldRole="")
modelerDataModel<-data.frame(modelerDataModel,var1)""")
```

表 68. *extensionprocessnode* プロパティ

extensionprocessnode プロパティ	データ・タイプ	プロパティの説明
syntax_type	R Python	R または Python のどちらのスクリプトを実行するか指定します (R がデフォルトです)。
r_syntax	string	実行する R スクリプト・シンタックス。
python_syntax	string	実行する Python スクリプト・シンタックス。
use_batch_size	フラグ	バッチ処理を使用可能にします。
batch_size	integer	各バッチに含めるデータ レコードの数を指定します。
convert_flags	StringsAndDoubles LogicalValues	フラグ型フィールドを変換するためのオプション。
convert_missing	フラグ	欠損値を R の NA 値に変換するオプションです。
convert_datetime	フラグ	日付形式または日付/時刻形式の変数を R の日付/時刻形式に変換するためのオプション。

表 68. *extensionprocessnode* プロパティ (続き)

extensionprocessnode プロパティ	データ・タイプ	プロパティの説明
convert_datetime_class	POSIXct POSIXlt	日付形式または日付/時刻形式の変数のうち、どの形式の変数を変換するかを指定するためのオプション。

mergenode プロパティ



レコード結合ノードは、複数の入力レコードを取得し、入力フィールドの全部または一部を含む1つの出力レコードを作成します。この機能は、内部顧客データと購入人口データのような、異なるソースからのデータを結合する場合に役立ちます。

例

```
node = stream.create("merge", "My node")
# assume customerdata and salesdata are configured database import nodes
stream.link(customerdata, node)
stream.link(salesdata, node)
node.setPropertyValue("method", "Keys")
node.setPropertyValue("key_fields", ["id"])
node.setPropertyValue("common_keys", True)
node.setPropertyValue("join", "PartialOuter")
node.setKeyedPropertyValue("outer_join_tag", "2", True)
node.setKeyedPropertyValue("outer_join_tag", "4", True)
node.setPropertyValue("single_large_input", True)
node.setPropertyValue("single_large_input_tag", "2")
node.setPropertyValue("use_existing_sort_keys", True)
node.setPropertyValue("existing_sort_keys", [["id", "Ascending"]])
```

表 69. *mergenode* プロパティ

mergenode プロパティ	データ・タイプ	プロパティの説明
method	Order Keys Condition Rankedcondition	データ ファイルでのリスト順にレコードを結合するかどうか、1つ以上のキー フィールドを使用してキー フィールド内の同じ値にレコードを結合するかどうか、指定された条件を満たす場合にレコードを結合するかどうか、1次データセットとすべての2次データセット内の各行のペアを結合するかどうかを指定します。いずれの場合も、ランク付け式を使用して、ランクの低い一致からランクの高い一致の順にすべての一致がソートされます。
condition	string	method が Condition に設定されている場合、レコードを含めるまたは破棄する条件を指定します。
key_fields	リスト	
common_keys	フラグ	

表 69. *mergenode* プロパティ (続き)

mergenode プロパティ	データ・タイプ	プロパティの説明
join	Inner FullOuter PartialOuter Anti	
outer_join_tag.n	フラグ	このプロパティでは、 <i>n</i> は「データセットの選択」ダイアログ・ボックスに表示されるタグ名です。どのようなデータセット数であっても不完全なレコードを作成する可能性があるため、複数のタグ名を指定できます。
single_large_input	フラグ	ほかの入力と比べて比較的大きな入力を指定し最適化を行うかどうかを指定します。
single_large_input_tag	string	「ラージ・データセットの選択」ダイアログ・ボックスに表示されるタグ名を指定します。このプロパティの用途は、1つの入力データセットしか指定できないという点で、outer_join_tag プロパティとは若干異なることに注意してください (データ型がフラグと文字列という違いもあり)。
use_existing_sort_keys	フラグ	入力がすでにキー・フィールドでソート済みかどうかを指定します。
existing_sort_keys	[['文字列', 'Ascending'] ¥ ['文字列', 'Descending']]	すでにソートされたフィールドとソート方向を指定します。
primary_dataset	string	method が Rankedcondition の場合は、結合内の 1 次データセットを選択します。これは外部結合の左側と考えることができます。
rename_duplicate_fields	Boolean	method が Rankedcondition の場合にこのプロパティを Y に設定し、異なるデータソースから取得された同じ名前を持つ複数のフィールドが結果の結合データセットに含まれている場合、それらのデータソースの各タグがフィールドの列見出しの先頭に追加されます。
merge_condition	string	
ranking_expression	string	
Num_matches	integer	merge_condition と ranking_expression に基づいて返される一致の数。最小値は 1、最大値は 100 です。

rfmaggregatenode プロパティ



リーセンシ、フリクエンシ、マネタリー (RFM) 集計ノードを使用すると、顧客の過去のトランザクション・データを取得し、未使用のデータをすべて削除し、残りのすべてのトランザクション・データを単一行に結合して、ユーザーが最後にいつ処理したか、行ったトランザクションの数、およびそれらのトランザクションの合計金額をリストすることができます。

例

```
node = stream.create("rfmaggregate", "My node")
node.setPropertyValue("relative_to", "Fixed")
node.setPropertyValue("reference_date", "2007-10-12")
node.setPropertyValue("id_field", "CardID")
node.setPropertyValue("date_field", "Date")
node.setPropertyValue("value_field", "Amount")
node.setPropertyValue("only_recent_transactions", True)
node.setPropertyValue("transaction_date_after", "2000-10-01")
```

表 70. rfmaggregatenode プロパティ

rfmaggregatenode プロパティ	データ・タイプ	プロパティの説明
relative_to	Fixed Today	トランザクションのリーセンシが計算される日付を指定します。
reference_date	日付	Fixed が relative_to に設定されている場合にのみ使用できます。
contiguous	フラグ	データ・ストリーム中で同じ ID を持つすべてのレコードが一緒に表示されるようにデータをソートしている場合、このオプションを選択すると処理を高速化することができます。
id_field	フィールド	顧客およびトランザクションを識別するために使用するフィールドを指定します。
date_field	フィールド	リーセンシを計算するために使用される日付フィールドを選択します。
value_field	フィールド	マネタリー値を計算するために使用するフィールドを指定します。
extension	string	重複集計フィールドに対応させる接頭辞または接尾辞を指定します。
add_as	Suffix Prefix	extension を接尾辞として追加するか、または接頭辞として追加するかを指定します。
discard_low_value_records	フラグ	discard_records_below 設定の使用を有効にします。
discard_records_below	number	RFM の合計を計算する場合に使用されないトランザクションの詳細の最小値を指定することができます。値の単位は、選択した value フィールドに関連します。

表 70. rfmaggregatenode プロパティ (続き)

rfmaggregatenode プロパティ	データ・タイプ	プロパティの説明
only_recent_transactions	フラグ	specify_transaction_date または transaction_within_last 設定の使用を有効にします。
specify_transaction_date	フラグ	
transaction_date_after	日付	specify_transaction_date が選択されている場合にのみ使用できます。データが分析に含まれた後のトランザクションの日付を指定します。
transaction_within_last	number	transaction_within_last が選択されている場合にのみ使用できます。レコードが分析に含まれる後の「リーセンシ基準日」の日付からさかのぼった期間の数および種類 (日、週、月または年数) を指定します。
transaction_scale	Days Weeks Months Years	transaction_within_last が選択されている場合にのみ使用できます。レコードが分析に含まれる後の「リーセンシ基準日」の日付からさかのぼった期間の数および種類 (日、週、月または年数) を指定します。
save_r2	フラグ	各顧客の 2 番目に最近のトランザクションの日付を表示します。
save_r3	フラグ	save_r2 が選択されている場合にのみ使用できます。各顧客の 3 番目に最近のトランザクションの日付を表示します。

Rprocessnode プロパティ



R 変換ノードでは、IBM(r) SPSS(r) Modeler ストリームからデータを取得し、そのデータを独自のカスタム R スクリプトを使用して変更できます。データ変更後、データはストリームに返されます。

例

```
node = stream.create("rprocess", "My node")
node.setPropertyValue("custom_name", "my_node")
node.setPropertyValue("syntax", """"day<-as.Date(modelerData$dob, format="%Y-%m-%d")
next_day<-day + 1
modelerData<-cbind(modelerData,next_day)
var1<-c(fieldName="Next
day",fieldLabel="",fieldStorage="date",fieldMeasure="",fieldFormat="",
fieldRole="")
modelerDataModel<-data.frame(modelerDataModel,var1)""")
node.setPropertyValue("convert_datetime", "POSIXct")
```


表 71. Rprocessnode プロパティ

Rprocessnode プロパティ	データ型	プロパティの説明
syntax	string	
convert_flags	StringsAndDoubles LogicalValues	
convert_datetime	フラグ	
convert_datetime_class	POSIXct POSIXlt	
convert_missing	フラグ	
use_batch_size	フラグ	バッチ処理を使用可能にします
batch_size	integer	各バッチに含めるデータ レコードの数を指定します

samplenode プロパティ



サンプリング・ノードでは、レコードのサブセットを選択します。層化サンプル、クラスター・サンプル、非無作為 (構造化) サンプルなど、さまざまなサンプルの種類がサポートされています。サンプリングは、パフォーマンスの向上、および分析のための関連するレコードまたはトランザクションのグループの選択に役に立ちます。

例

```
/* Create two Sample nodes to extract
different samples from the same data */

node = stream.create("sample", "My node")
node.setPropertyValue("method", "Simple")
node.setPropertyValue("mode", "Include")
node.setPropertyValue("sample_type", "First")
node.setPropertyValue("first_n", 500)

node = stream.create("sample", "My node")
node.setPropertyValue("method", "Complex")
node.setPropertyValue("stratify_by", ["Sex", "Cholesterol"])
node.setPropertyValue("sample_units", "Proportions")
node.setPropertyValue("sample_size_proportions", "Custom")
node.setPropertyValue("sizes_proportions", [["M", "High", "Default"], ["M",
"Normal", "Default"],
["F", "High", 0.3], ["F", "Normal", 0.3]])
```

表 72. samplenode プロパティ

samplenode プロパティ	データ・タイプ	プロパティの説明
method	単純 Complex	
mode	Include Discard	指定された条件を満たすレコードを含めるか (Include)、破棄 (Discard) します。

表 72. *samplenode* プロパティ (続き)

samplenode プロパティ	データ・タイプ	プロパティの説明
sample_type	First OneInN RandomPct	サンプリング方法を指定します。
first_n	integer	指定された分割点までのレコードを含めるか破棄します。
one_in_n	number	<i>n</i> 番目ごとにレコードを含めるか破棄します。
rand_pct	number	含めるか破棄するレコードのパーセンテージを指定します。
use_max_size	フラグ	maximum_size 設定の使用を有効にします。
maximum_size	integer	データ・ストリームに入れるまたはデータ・ストリームから破棄するサンプルの最大数を指定します。このオプションは冗長であり、そのため、First と Include が指定されているときは破棄されます。
set_random_seed	フラグ	ランダム・シード設定の使用を有効にします。
random_seed	integer	ランダム・シードとして使用する値を指定します。
complex_sample_type	Random Systematic	
sample_units	Proportions Counts	
sample_size_proportions	Fixed Custom Variable	
sample_size_counts	Fixed Custom Variable	
fixed_proportions	number	
fixed_counts	integer	
variable_proportions	フィールド	
variable_counts	フィールド	
use_min_stratum_size	フラグ	

表 72. <i>samplenode</i> プロパティ (続き)		
samplenode プロパティ	データ・タイプ	プロパティの説明
minimum_stratum_size	<i>integer</i>	このオプションは、Sample units=Proportions でコンプレックス・サンプルが取得される場合にのみ適用されます。
use_max_stratum_size	フラグ	
maximum_stratum_size	<i>integer</i>	このオプションは、Sample units=Proportions でコンプレックス・サンプルが取得される場合にのみ適用されます。
clusters	フィールド	
stratify_by	[<i>field1 ... fieldN</i>]	
specify_input_weight	フラグ	
input_weight	フィールド	
new_output_weight	<i>string</i>	
sizes_proportions	[[<i>string string value</i>] [<i>string string value</i>]...]	sample_units=proportions および sample_size_proportions=Custom の場合、層化フィールドの値の考えられる組み合わせの値を指定します。
default_proportion	<i>number</i>	
sizes_counts	[[<i>string string value</i>] [<i>string string value</i>]...]	層化フィールドの値の考えられる組み合わせの値を指定します。使用 방법은 sizes_proportions と似ていますが、割合ではなく整数を指定します。
default_count	<i>number</i>	

selectnode プロパティ



条件抽出ノードで、特定の条件に基づいて、データ・ストリームからレコードのサブセットを選択したり破棄したりできます。例えば、特定の営業地域に関連するレコードを選択できます。

例

```
node = stream.create("select", "My node")
node.setPropertyValue("mode", "Include")
node.setPropertyValue("condition", "Age < 18")
```

表 73. <i>selectnode</i> プロパティ		
selectnode プロパティ	データ・タイプ	プロパティの説明
mode	Include	選択したレコードを含めるか、または破棄するかを指定します。
	Discard	
condition	<i>string</i>	レコードを含めるか、または破棄かの条件。

sortnode プロパティ



ソート・ノードで、1つまたは複数のフィールド値に基づいて、レコードを昇順または降順にソートします。

例

```
node = stream.create("sort", "My node")
node.setPropertyValue("keys", [["Age", "Ascending"], ["Sex", "Descending"]])
node.setPropertyValue("default_ascending", False)
node.setPropertyValue("use_existing_keys", True)
node.setPropertyValue("existing_keys", [["Age", "Ascending"]])
```

表 74. sortnode プロパティ

sortnode プロパティ	データ・タイプ	プロパティの説明
keys	リスト	ソートの基準となるフィールドを指定します。ソートの方向が指定されていない場合、デフォルトが使用されます。
default_ascending	フラグ	デフォルトのソート順を指定します。
use_existing_keys	フラグ	前に使用されたフィールドのソート順を使用してソートを最適化するかどうかを指定します。
existing_keys		すでにソートされたフィールドとソート方向を指定します。keys プロパティと同じ形式を使用します。

spacetimeboxes プロパティ



スペース・タイム・ボックス (STB) は、Geohash の空間的な場所を拡張したものです。具体的には、STB は英数字の文字列で、空間および時間を規則的に分割した領域です。

表 75. spacetimeboxes プロパティ

spacetimeboxes プロパティ	データ・タイプ	プロパティの説明
mode	<i>IndividualRecords</i> <i>Hangouts</i>	
latitude_field	フィールド	
longitude_field	フィールド	
timestamp_field	フィールド	

表 75. *spacetimeboxes* プロパティ (続き)

spacetimeboxes プロパティ	データ・タイプ	プロパティの説明
densities	[<i>density, density, density...</i>]	各 <i>density</i> は、文字列です。例: STB_GH8_1DAY <i>densities</i> が有効であるための制限が存在することに注意してください。 geohash の場合、GH1-GH15 の値を使用できます。 この部分では、以下の値を使用できます <pre> EVER 1YEAR 1MONTH 1DAY 12HOURS 8HOURS 6HOURS 4HOURS 3HOURS 2HOURS 1HOUR 30MINS 15MINS 10MINS 5MINS 2 MINS 1 MIN 30SECS 15SECS 10SECS 5 SECS 2 SECS 1SEC </pre>
field_name_extension	<i>string</i>	
add_extension_as	<i>prefix</i> <i>Suffix</i>	
hangout_density	密度	単一密度 (上記を参照)
id_field	フィールド	
qualifying_duration	<pre> 1DAY 12HOURS 8HOURS 6HOURS 4HOURS 2HOURS 1HOUR 30MIN 15MIN 10MIN 5MIN 2MIN 1MIN 30SECS 15SECS 10SECS 5SECS 2SECS 1SECS </pre>	これは、文字列にする必要があります。
min_events	<i>integer</i>	最小値は 2 です。
qualifying_pct	<i>integer</i>	1 から 100 の範囲にする必要があります

streamingtimeseries プロパティ



ストリーミング時系列分析ノードは、1つのステップで時系列モデルを作成してスコアリングします。

注: このストリーミング時系列分析ノードは、SPSS モデラー バージョン 18 で廃止されたオリジナルのストリーミング時系列分析ノードに置き換わるものです。

表 76. streamingtimeseries プロパティ		
streamingtimeseries プロパティ	値	プロパティの説明
targets	フィールド	ストリーミング時系列分析ノードは、オプションで1つ以上の入力フィールドを予測値として使用しながら、1つ以上の対象フィールドを予測します。度数フィールドおよび重みフィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
candidate_inputs	[field1 ... fieldN]	モデルで使用される入力または予測変数フィールド。
use_period	フラグ	
date_time_field	フィールド	

表 76. <i>streamingtimeseries</i> プロパティ (続き)		
streamingtimeseries プロパティ	値	プロパティの説明
input_interval	None	
	Unknown	
	Year	
	Quarter	
	Month	
	Week	
	Day	
	Hour	
	Hour_nonperiod	
	Minute	
	Minute_nonperiod	
	Second	
	Second_nonperiod	
period_field	フィールド	
period_start_value	<i>integer</i>	
num_days_per_week	<i>integer</i>	
start_day_of_week	Sunday	
	Monday	
	Tuesday	
	Wednesday	
	Thursday	
	Friday	
	Saturday	
num_hours_per_day	<i>integer</i>	
start_hour_of_day	<i>integer</i>	
timestamp_increments	<i>integer</i>	

表 76. *streamingtimeseries* プロパティ (続き)

streamingtimeseries プロパティ	値	プロパティの説明
cyclic_increments	<i>integer</i>	
cyclic_periods	リスト	
output_interval	None Year Quarter Month Week Day Hour Minute Second	
is_same_interval	フラグ	
cross_hour	フラグ	
aggregate_and_distribute	リスト	
aggregate_default	Mean Sum Mode Min Max	
distribute_default	Mean Sum	
group_default	Mean Sum Mode Min Max	

表 76. *streamingtimeseries* プロパティ (続き)

streamingtimeseries プロパティ	値	プロパティの説明
missing_imput	Linear_interp Series_mean K_mean K_median Linear_trend	
k_span_points	<i>integer</i>	
use_estimation_period	フラグ	
estimation_period	Observations Times	
date_estimation	リスト	<code>date_time_field</code> を使用する場合にのみ使用可能です
period_estimation	リスト	<code>use_period</code> を使用する場合にのみ使用可能です
observations_type	Latest Earliest	
observations_num	<i>integer</i>	
observations_exclude	<i>integer</i>	
method	ExpertModeler Exsmooth Arima	
expert_modeler_method	ExpertModeler Exsmooth Arima	
consider_seasonal	フラグ	
detect_outliers	フラグ	
expert_outlier_additive	フラグ	
expert_outlier_level_shift	フラグ	
expert_outlier_innovational	フラグ	
expert_outlier_level_shift	フラグ	

表 76. *streamingtimeseries* プロパティ (続き)

streamingtimeseries プロパティ	値	プロパティの説明
expert_outlier_transient	フラグ	
expert_outlier_seasonal_additive	フラグ	
expert_outlier_local_trend	フラグ	
expert_outlier_additive_patch	フラグ	
consider_newesmodels	フラグ	
exsmooth_model_type	Simple HoltsLinearTrend BrownsLinearTrend DampedTrend SimpleSeasonal WintersAdditive WintersMultiplicative DampedTrendAdditive DampedTrendMultiplicative MultiplicativeTrendAdditive MultiplicativeSeasonal MultiplicativeTrendMultiplicative MultiplicativeTrend	
futureValue_type_method	Compute specify	
exsmooth_transformation_type	None SquareRoot NaturalLog	
arima.p	integer	

表 76. *streamingtimeseries* プロパティ (続き)

streamingtimeseries プロパティ	値	プロパティの説明
<code>arma.d</code>	<i>integer</i>	
<code>arma.q</code>	<i>integer</i>	
<code>arma.sp</code>	<i>integer</i>	
<code>arma.sd</code>	<i>integer</i>	
<code>arma.sq</code>	<i>integer</i>	
<code>arma_transformation_type</code>	None SquareRoot NaturalLog	
<code>arma_include_constant</code>	フラグ	
<code>tf_arma.p.fieldname</code>	<i>integer</i>	伝達関数用。
<code>tf_arma.d.fieldname</code>	<i>integer</i>	伝達関数用。
<code>tf_arma.q.fieldname</code>	<i>integer</i>	伝達関数用。
<code>tf_arma.sp.fieldname</code>	<i>integer</i>	伝達関数用。
<code>tf_arma.sd.fieldname</code>	<i>integer</i>	伝達関数用。
<code>tf_arma.sq.fieldname</code>	<i>integer</i>	伝達関数用。
<code>tf_arma.delay.fieldname</code>	<i>integer</i>	伝達関数用。
<code>tf_arma.transformation_type.fieldname</code>	None SquareRoot NaturalLog	伝達関数用。
<code>arma_detect_outliers</code>	フラグ	
<code>arma_outlier_additive</code>	フラグ	
<code>arma_outlier_level_shift</code>	フラグ	
<code>arma_outlier_innovational</code>	フラグ	
<code>arma_outlier_transient</code>	フラグ	
<code>arma_outlier_seasonal_additive</code>	フラグ	
<code>arma_outlier_local_trend</code>	フラグ	
<code>arma_outlier_additive_patch</code>	フラグ	
<code>conf_limit_pct</code>	<i>real</i>	
<code>events</code>	フィールド	
<code>forecastperiods</code>	<i>integer</i>	
<code>extend_records_into_future</code>	フラグ	
<code>conf_limits</code>	フラグ	

表 76. *streamingtimeseries* プロパティ (続き)

streamingtimeseries プロパティ	値	プロパティの説明
noise_res	フラグ	

streamingts プロパティ (廃止)



注: この元のストリーミング時系列分析ノードは スポス モデラー のバージョン 18 で廃止され、新しいストリーミング時系列分析ノードに置き換えられました。この新しいノードは、IBMSPSS Analytic Server の機能を活用し、ビッグデータを処理するように設計されています。

ストリーミング TS ノードは、1つのステップで時系列モデルを作成してスコアリングします。時間間隔ノードは必要ありません。

例

```
node = stream.create("streamingts", "My node")
node.setPropertyValue("deployment_force_rebuild", True)
node.setPropertyValue("deployment_rebuild_mode", "Count")
node.setPropertyValue("deployment_rebuild_count", 3)
node.setPropertyValue("deployment_rebuild_pct", 11)
node.setPropertyValue("deployment_rebuild_field", "Year")
```

表 77. *streamingts* プロパティ

streamingts プロパティ	データ・タイプ	プロパティの説明
custom_fields	フラグ	custom_fields=false の場合は、上流のデータ型ノードの設定が使用されます。 custom_fields=true の場合は、targets と inputs を指定する必要があります。
targets	[field1...fieldN]	
inputs	[field1...fieldN]	
method	ExpertModeler Exsmooth Arima	
calculate_conf	フラグ	
conf_limit_pct	real	
use_time_intervals_node	フラグ	use_time_intervals_node=true の場合は、上流の時間区分ノードの設定が使用されます。 use_time_intervals_node=false の場合は、interval_offset_position、interval_offset、および interval_type を指定する必要があります。
interval_offset_position	LastObservation LastRecord	LastObservation は、「最新の有効な観測値」を表します。LastRecord は、「最後のレコードから遡り設定」を表します。
interval_offset	number	

表 77. *streamingts* プロパティ (続き)

streamingts プロパティ	データ・タイプ	プロパティの説明
interval_type	Periods Years Quarters Months WeeksNonPeriodic DaysNonPeriodic HoursNonPeriodic MinutesNonPeriodic SecondsNonPeriodic	
events	フィールド	
expert_modeler_method	AllModels Exsmooth Arima	
consider_seasonal	フラグ	
detect_outliers	フラグ	
expert_outlier_additive	フラグ	
expert_outlier_level_shift	フラグ	
expert_outlier_innovational	フラグ	
expert_outlier_transient	フラグ	
expert_outlier_seasonal_additive	フラグ	
expert_outlier_local_trend	フラグ	
expert_outlier_additive_patch	フラグ	
exsmooth_model_type	Simple HoltsLinearTrend BrownsLinearTrend DampedTrend SimpleSeasonal WintersAdditive WintersMultiplicative	
exsmooth_transformation_type	None SquareRoot NaturalLog	
arima_p	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ
arima_d	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ
arima_q	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ

表 77. *streamingts* プロパティ (続き)

streamingts プロパティ	データ・タイプ	プロパティの説明
<code>arma_sp</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ
<code>arma_sd</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ
<code>arma_sq</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ
<code>arma_transformation_type</code>	None SquareRoot NaturalLog	時系列モデル作成ノードの場合と同じプロパティ
<code>arma_include_constant</code>	フラグ	時系列モデル作成ノードの場合と同じプロパティ
<code>tf_arma_p.fieldname</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ。 伝達関数用。
<code>tf_arma_d.fieldname</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ。 伝達関数用。
<code>tf_arma_q.fieldname</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ。 伝達関数用。
<code>tf_arma_sp.fieldname</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ。 伝達関数用。
<code>tf_arma_sd.fieldname</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ。 伝達関数用。
<code>tf_arma_sq.fieldname</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ。 伝達関数用。
<code>tf_arma_delay.fieldname</code>	<i>integer</i>	時系列モデル作成ノードの場合と同じプロパティ。 伝達関数用。
<code>tf_arma_transformation_type.fieldname</code>	None SquareRoot NaturalLog	
<code>arma_detect_outlier_mode</code>	None Automatic	
<code>arma_outlier_additive</code>	フラグ	
<code>arma_outlier_level_shift</code>	フラグ	
<code>arma_outlier_innovational</code>	フラグ	
<code>arma_outlier_transient</code>	フラグ	
<code>arma_outlier_seasonal_additive</code>	フラグ	
<code>arma_outlier_local_trend</code>	フラグ	

表 77. <i>streamingts</i> プロパティ (続き)		
streamingts プロパティ	データ・タイプ	プロパティの説明
arima_outlier_additive_patch	フラグ	
deployment_force_rebuild	フラグ	
deployment_rebuild_mode	Count Percent	
deployment_rebuild_count	<i>number</i>	
deployment_rebuild_pct	<i>number</i>	
deployment_rebuild_field	< <i>field</i> >	

第 11 章 フィールド設定ノードのプロパティ

anonymizenode プロパティ



匿名化ノードは、フィールド名や値の下流の表示方法を変換し、元のデータを隠します。これは、他のユーザーが顧客名やその他の詳細情報をなどの重要情報を使用してモデルを構築できるようにする場合に有用です。

例

```
stream = modeler.script.stream()
varfilenode = stream.createAt("variablefile", "File", 96, 96)
varfilenode.setPropertyValue("full_filename", "$CLEO/DEMOS/DRUG1n")
node = stream.createAt("anonymize", "My node", 192, 96)
# Anonymize node requires the input fields while setting the values
stream.link(varfilenode, node)
node.setKeyedPropertyValue("enable_anonymize", "Age", True)
node.setKeyedPropertyValue("transformation", "Age", "Random")
node.setKeyedPropertyValue("set_random_seed", "Age", True)
node.setKeyedPropertyValue("random_seed", "Age", 123)
node.setKeyedPropertyValue("enable_anonymize", "Drug", True)
node.setKeyedPropertyValue("use_prefix", "Drug", True)
node.setKeyedPropertyValue("prefix", "Drug", "myprefix")
```

表 78. anonymizenode プロパティ

anonymizenode プロパティ	データ・タイプ	プロパティの説明
enable_anonymize	フラグ	これを True に設定すると、フィールド値の匿名化がアクティブになります (そのフィールドの「匿名値」列で「はい」を選択した場合も、同じ結果になります)。
use_prefix	フラグ	これを True に設定すると、ユーザー指定の接頭辞が使用されます (ユーザー指定の接頭辞が指定されている場合)。ハッシュ・メソッドによって匿名化されるフィールドに適用され、そのフィールドの「値を置換」ダイアログの「ユーザー設定」ラジオ・ボタンを選択することと同等です。
prefix	string	「値を置換」ダイアログ・ボックスのテキスト・ボックスに接頭辞を入力することと同等です。デフォルトの接頭辞は、何も他に指定されていない場合は、デフォルト値です。
transformation	Random Fixed	Transform メソッドにより匿名化されたフィールドの変換パラメーターが無作為 (Random) か固定 (Fixed) かを決定します。
set_random_seed	フラグ	これを True に設定すると、指定されたシード値が使用されます (transformation も Random に設定されている場合)。
random_seed	integer	set_random_seed が True に設定されている場合、このプロパティは乱数のシードになります。

表 78. anonymizenode プロパティ (続き)

anonymizenode プロパティ	データ・タイプ	プロパティの説明
scale	number	transformation が Fixed に設定されている場合、この値はスケール用として使用されます。最大スケール値は通常 10 ですが、あふれを防止するために減少できます。
translate	number	transformation が Fixed に設定されている場合、この値は変換用として使用されます。最大変換値は通常 1000 ですが、あふれを防止するために減少できます。

autodatapreprenode プロパティ



データの自動準備 (ADP) ノードでは、データ分析、固定値の識別、問題のあるまたは役に立たない可能性のあるフィールドのスクリーニング、必要に応じた新しい属性の取得、詳細なスクリーニングおよびサンプリング手法を使用したパフォーマンスの向上などを行うことができます。完全に自動化された方法でノードを使用し、ノードで固定値を選択および適用できます。または必要に応じて変更の作成および承認、拒否または修正の前に変更をプレビューできます。

例

```
node = stream.create("autodataprep", "My node")
node.setPropertyValue("objective", "Balanced")
node.setPropertyValue("excluded_fields", "Filter")
node.setPropertyValue("prepare_dates_and_times", True)
node.setPropertyValue("compute_time_until_date", True)
node.setPropertyValue("reference_date", "Today")
node.setPropertyValue("units_for_date_durations", "Automatic")
```

表 79. autodatapreprenode プロパティ

autodatapreprenode プロパティ	データ・タイプ	プロパティの説明
objective	Balanced Speed Accuracy Custom	
custom_fields	フラグ	真 (true) の場合は、現在のノードのターゲット、入力、その他フィールドなどを指定することができます。偽 (false) の場合は、上流のデータ型ノードから現在の設定が使用されます。
target	フィールド	1つの対象フィールドを指定します。
inputs	[field1 ... fieldN]	モデルで使用する入力または予測変数フィールド。
use_frequency	フラグ	
frequency_field	フィールド	

表 79. autodatapreprenode プロパティ (続き)

autodatapreprenode プロパティ	データ・タイプ	プロパティの説明
use_weight	フラグ	
weight_field	フィールド	
excluded_fields	Filter None	
if_fields_do_not_match	StopExecution ClearAnalysis	
prepare_dates_and_times	フラグ	すべての日付/時間フィールドへのアクセスを制御します。
compute_time_until_date	フラグ	
reference_date	Today Fixed	
fixed_date	日付	
units_for_date_durations	Automatic Fixed	
fixed_date_units	Years Months Days	
compute_time_until_time	フラグ	
reference_time	CurrentTime Fixed	
fixed_time	時間	
units_for_time_durations	Automatic Fixed	
fixed_date_units	Hours Minutes Seconds	
extract_year_from_date	フラグ	
extract_month_from_date	フラグ	
extract_day_from_date	フラグ	

表 79. autodatapreprenode プロパティ (続き)

autodatapreprenode プロパティ	データ・タイプ	プロパティの説明
extract_hour_from_time	フラグ	
extract_minute_from_time	フラグ	
extract_second_from_time	フラグ	
exclude_low_quality_inputs	フラグ	
exclude_too_many_missing	フラグ	
maximum_percentage_missing	number	
exclude_too_many_categories	フラグ	
maximum_number_categories	number	
exclude_if_large_category	フラグ	
maximum_percentage_category	number	
prepare_inputs_and_target	フラグ	
adjust_type_inputs	フラグ	
adjust_type_target	フラグ	
reorder_nominal_inputs	フラグ	
reorder_nominal_target	フラグ	
replace_outliers_inputs	フラグ	
replace_outliers_target	フラグ	
replace_missing_continuous_inputs	フラグ	
replace_missing_continuous_target	フラグ	
replace_missing_nominal_inputs	フラグ	
replace_missing_nominal_target	フラグ	
replace_missing_ordinal_inputs	フラグ	
replace_missing_ordinal_target	フラグ	
maximum_values_for_ordinal	number	
minimum_values_for_continuous	number	
outlier_cutoff_value	number	

表 79. autodatapreprenode プロパティ (続き)

autodatapreprenode プロパティ	データ・タイプ	プロパティの説明
outlier_method	Replace Delete	
rescale_continuous_inputs	フラグ	
rescaling_method	MinMax ZScore	
min_max_minimum	number	
min_max_maximum	number	
z_score_final_mean	number	
z_score_final_sd	number	
rescale_continuous_target	フラグ	
target_final_mean	number	
target_final_sd	number	
transform_select_input_fields	フラグ	
maximize_association_with_target	フラグ	
p_value_for_merging	number	
merge_ordinal_features	フラグ	
merge_nominal_features	フラグ	
minimum_cases_in_category	number	
bin_continuous_fields	フラグ	
p_value_for_binning	number	
perform_feature_selection	フラグ	
p_value_for_selection	number	
perform_feature_construction	フラグ	
transformed_target_name_extension	string	
transformed_inputs_name_extension	string	
constructed_features_root_name	string	
years_duration_name_extension	string	

表 79. autodatapreprenode プロパティ (続き)

autodatapreprenode プロパティ	データ・タイプ	プロパティの説明
months_duration_name_extension	string	
days_duration_name_extension	string	
hours_duration_name_extension	string	
minutes_duration_name_extension	string	
seconds_duration_name_extension	string	
year_cyclical_name_extension	string	
month_cyclical_name_extension	string	
day_cyclical_name_extension	string	
hour_cyclical_name_extension	string	
minute_cyclical_name_extension	string	
second_cyclical_name_extension	string	

astimeintervalsnode プロパティ



間隔を指定し、新しい時間フィールドを作成して推定や予測を行う場合は、時間区分ノードを使用します。全時間区分がサポートされ、範囲は秒から数年の間を指定することができます。

表 80. astimeintervalsnode プロパティ

astimeintervalsnode プロパティ	データ・タイプ	プロパティの説明
time_field	フィールド	1つの連続型フィールドのみ許可されます。ノードは、このフィールドを集計キーとして使用して、間隔を変換します。ここで整数フィールドを使用すると、そのフィールドは時間インデックスとして認識されます。
dimensions	[field1 field2 ... fieldn]	これらのフィールドを使用して、各フィールドの値に基づき、個々の時系列が作成されます。

表 80. *astimeintervalnode* プロパティ (続き)

astimeintervalnode プロパティ	データ・タイプ	プロパティの説明
fields_to_aggregate	<i>[field1 field2 ... fieldn]</i>	これらのフィールドは、時間フィールドの期間変更処理の一部として集計されます。このピッカーに含まれていないすべてのフィールドが、ノードから送信されるデータから除外されます。

binningnode プロパティ



データ分割ノードで、既存の 1 つまたは複数の連続型 (数値範囲) フィールドの値に基づいて、自動的に新しい名義型 (セット型) フィールドを作成します。例えば、連続型収入フィールドを、平均からの偏差による収入グループを含む、新しいカテゴリー・フィールドに変換することができます。新規フィールドのビンを作成すると、分割点に基づいてフィールド作成ノードを生成することができます。

例

```
node = stream.create("binning", "My node")
node.setPropertyValue("fields", ["Na", "K"])
node.setPropertyValue("method", "Rank")
node.setPropertyValue("fixed_width_name_extension", "_binned")
node.setPropertyValue("fixed_width_add_as", "Suffix")
node.setPropertyValue("fixed_bin_method", "Count")
node.setPropertyValue("fixed_bin_count", 10)
node.setPropertyValue("fixed_bin_width", 3.5)
node.setPropertyValue("tile10", True)
```

表 81. *binningnode* プロパティ

binningnode プロパティ	データ・タイプ	プロパティの説明
fields	<i>[field1 field2 ... fieldn]</i>	変換保留中の連続型 (数値範囲) フィールド。複数のフィールドを同時にビンに分割できます。
method	FixedWidth EqualCount Rank SDev Optimal	新規フィールドのビン (カテゴリー) の分割点を決める方法。
recalculate_bins	Always IfNecessary	ノードが実行されるごとに、ビンが再計算され、適切なビンの中にデータが配置されるか、またはデータが既存のビンおよび追加された新規のビンに追加されるだけかを指定します。
fixed_width_name_extension	string	デフォルトの拡張子は <code>_BIN</code> です。

表 81. binningnode プロパティ (続き)

binningnode プロパティ	データ・タイプ	プロパティの説明
fixed_width_add_as	Suffix Prefix	拡張子をフィールド名の最後に追加するか (Suffix)、または先頭に追加するか (Prefix) を指定します。デフォルトの拡張子は <i>income_BIN</i> です。
fixed_bin_method	Width Count	
fixed_bin_count	integer	新規フィールドの固定幅ビン (カテゴリ) 数を決定するのに使用する整数を指定します。
fixed_bin_width	real	ビンの幅を算出するために使用する値 (整数または実数)。
equal_count_name_ extension	string	デフォルトの拡張子は <i>_TILE</i> です。
equal_count_add_as	Suffix Prefix	標準の分位を使用して生成されるフィールドに対して使用される拡張子が、Suffix (接頭辞) か Prefix (接尾辞) かを指定します。デフォルトの拡張子は、 <i>_TILE</i> に <i>N</i> を付けたものになります。 <i>N</i> は分位数です。
tile4	フラグ	それぞれが 25 % のケースを含む、4 分位のビンを生成します。
tile5	フラグ	5 つの 5 分位ビンを生成します。
tile10	フラグ	10 個の十分位 (デシル) ビンを生成します。
tile20	フラグ	20 個の二十分位ビンを生成します。
tile100	フラグ	100 個の百分位 (パーセンタイル) ビンを生成します。
use_custom_tile	フラグ	
custom_tile_name_extension	string	デフォルトの拡張子は <i>_TILEN</i> です。
custom_tile_add_as	Suffix Prefix	
custom_tile	integer	
equal_count_method	RecordCount ValueSum	RecordCount の方法は、同じ数のレコードを各ビンに割り当てます。一方、ValueSum では、各ビンの値の合計が同じになるようにレコードを割り当てます。
tied_values_method	Next Current Random	可否同数の値のデータに配置されるビン指定。

表 81. binningnode プロパティ (続き)		
binningnode プロパティ	データ・タイプ	プロパティの説明
rank_order	Ascending Descending	このプロパティには、Ascending (もっとも小さい値が 1 となる) または Descending (もっとも大きい値が 1 となる) が含まれます。
rank_add_as	Suffix Prefix	このオプションは、ランク、ランクの比率、およびランクのパーセンテージに適用されます。
rank	フラグ	
rank_name_extension	string	デフォルトの拡張子は <code>_RANK</code> です。
rank_fractional	フラグ	新規フィールドの値が、ランクを非欠損ケースの重みの合計で除算した値になるように、ケースをランク付けします。小数点付き順位の範囲は 0 から 1 までです。
rank_fractional_name_ extension	string	デフォルトの拡張子は <code>_F_RANK</code> です。
rank_pct	フラグ	各ランクが、有効な値を持つレコード数で除算された後、100 倍されます。パーセンテージの小数点付き順位の範囲は 1 から 100 までです。
rank_pct_name_extension	string	デフォルトの拡張子は <code>_P_RANK</code> です。
sdev_name_extension	string	
sdev_add_as	Suffix Prefix	
sdev_count	One Two Three	
optimal_name_extension	string	デフォルトの拡張子は <code>_OPTIMAL</code> です。
optimal_add_as	Suffix Prefix	
optimal_supervisor_field	フィールド	データ分割のために選択されたフィールドが関係する監督フィールドとして選ばれたフィールド。
optimal_merge_bins	フラグ	ケース度数が小さいビンをより大きな近傍ビンに追加することを指定します。
optimal_small_bin_threshold	integer	
optimal_pre_bin	フラグ	データセットの事前データ分割を実行することを示します。

表 81. binningnode プロパティ (続き)

binningnode プロパティ	データ・タイプ	プロパティの説明
optimal_max_bins	integer	過度に多数のビンを作成しないように、上限を指定します。
optimal_lower_end_point	Inclusive	
	Exclusive	
optimal_first_bin	Unbounded	
	Bounded	
optimal_last_bin	Unbounded	
	Bounded	

derivenode プロパティ



フィールド作成ノードで、1つまたは複数の既存フィールドから、データ値を変更するか、新しいフィールドを作成します。これで、種類が式、フラグ型、名義、状態、カウント、および条件式の各フィールドが作成されます。

例 1

```
# Create and configure a Flag Derive field node
node = stream.create("derive", "My node")
node.setPropertyValue("new_name", "DrugX_Flag")
node.setPropertyValue("result_type", "Flag")
node.setPropertyValue("flag_true", "1")
node.setPropertyValue("flag_false", "0")
node.setPropertyValue("flag_expr", "'Drug' == \"drugX\"")

# Create and configure a Conditional Derive field node
node = stream.create("derive", "My node")
node.setPropertyValue("result_type", "Conditional")
node.setPropertyValue("cond_if_cond", "@OFFSET(\"Age\", 1) = \"Age\"")
node.setPropertyValue("cond_then_expr", "(@OFFSET(\"Age\", 1) = \"Age\" >< @INDEX)")
node.setPropertyValue("cond_else_expr", "\"Age\"")
```

例 2

このスクリプトは、特定のポイント (特定のイベントが発生した場所など) の X 座標と Y 座標を表す XPos と YPos という 2 つの数値列があることを前提としています。このスクリプトにより、特定の座標系でその点を表す X 座標と Y 座標から地理空間列を計算するフィールド作成ノードが作成されます。

```
stream = modeler.script.stream()
# Other stream configuration code
node = stream.createAt("derive", "Location", 192, 96)
node.setPropertyValue("new_name", "Location")
node.setPropertyValue("formula_expr", "['XPos', 'YPos']")
node.setPropertyValue("formula_type", "Geospatial")
# Now we have set the general measurement type, define the
# specifics of the geospatial object
node.setPropertyValue("geo_type", "Point")
```

```
node.setPropertyValue("has_coordinate_system", True)
node.setPropertyValue("coordinate_system", "ETRS_1989_EPSG_Arctic_zone_5-47")
```

表 82. *derivenode* プロパティ

derivenode プロパティ	データ・タイプ	プロパティの説明
new_name	string	新しいフィールド名。
mode	Single Multiple	1つのフィールドか (Single)、または複数フィールドか (Multiple) を指定します。
fields	リスト	複数フィールドを選択する場合にだけ、Multiple モードで使用。
name_extension	string	新しいフィールド名に使用する拡張子を指定します。
add_as	Suffix Prefix	拡張子をフィールド名の Prefix (先頭、接頭辞)、または Suffix (最後、接尾辞) として追加します。
result_type	Formula Flag Set State Count Conditional	作成可能な新しいフィールドの 6 つの種類。
formula_expr	string	フィールド作成ノードの新しいフィールド値を計算する式。
flag_expr	string	
flag_true	string	
flag_false	string	
set_default	string	
set_value_cond	string	特定の値に関連付けられた条件を提供するように構造化プロパティ。
state_on_val	string	オン (On) の条件を満たす場合の新規フィールドの値を指定します。
state_off_val	string	オフ (Off) の条件を満たす場合の新規フィールドの値を指定します。
state_on_expression	string	
state_off_expression	string	

表 82. *derivenode* プロパティ (続き)

derivenode プロパティ	データ・タイプ	プロパティの説明
state_initial	On Off	各レコードで新しいフィールドの初期値として On または Off を割り当てます。この値は、それぞれの条件が満たされるごとに変化します。
count_initial_val	string	
count_inc_condition	string	
count_inc_expression	string	
count_reset_condition	string	
cond_if_cond	string	
cond_then_expr	string	
cond_else_expr	string	
formula_measure_type	Range / MeasureType.RANGE Discrete / MeasureType.DISCRETE Flag / MeasureType.FLAG Set / MeasureType.SET OrderedSet / MeasureType.ORDERED_SET Typeless / MeasureType.TYPELESS Collection / MeasureType.COLLECTION Geospatial / MeasureType.GEOSPATIAL	このプロパティを使用して、作成されたフィールドに関連付けられた尺度を定義することができます。setter 関数には、文字列か、MeasureType の値のいずれかを渡すことができます。getter は、常に MeasureType の値を返します。
collection_measure	Range / MeasureType.RANGE Flag / MeasureType.FLAG Set / MeasureType.SET OrderedSet / MeasureType.ORDERED_SET Typeless / MeasureType.TYPELESS	収集フィールド (深さが 0 のリスト) の場合、このプロパティは、基礎となる値に関連付けられた尺度タイプを定義します。

表 82. *derivenode* プロパティ (続き)

derivenode プロパティ	データ・タイプ	プロパティの説明
geo_type	Point MultiPoint LineString MultiLineString Polygon MultiPolygon	地理空間フィールドの場合、このプロパティにより、このフィールドが表す地理空間オブジェクトのタイプが定義されます。これは、値のリストの深さと整合している必要があります。
has_coordinate_system	Boolean	地理空間フィールドの場合、このプロパティにより、このフィールドに座標系があるかどうか定義されます。
coordinate_system	string	地理空間フィールドの場合、このプロパティにより、このフィールドの座標系が定義されます。

ensembledenode プロパティ



アンサンブル・ノードでは、2つ以上のモデル・ナゲットを組み合わせることで1つのモデルよりもより正確な予測を取得します。

例

```
# Create and configure an Ensemble node
# Use this node with the models in demos\streams\pm_binaryclassifier.stream
node = stream.create("ensemble", "My node")
node.setPropertyValue("ensemble_target_field", "response")
node.setPropertyValue("filter_individual_model_output", False)
node.setPropertyValue("flag_ensemble_method", "ConfidenceWeightedVoting")
node.setPropertyValue("flag_voting_tie_selection", "HighestConfidence")
```

表 83. *ensembledenode* プロパティ

ensembledenode プロパティ	データ・タイプ	プロパティの説明
ensemble_target_field	フィールド	アンサンブルで使用するすべてのモデルの対象フィールドを指定します。
filter_individual_model_output	フラグ	個々のモデルのスコアリング結果を抑制するかどうかを指定します。

表 83. *ensembledenode* プロパティ (続き)

ensembledenode プロパティ	データ・タイプ	プロパティの説明
flag_ensemble_method	Voting ConfidenceWeightedVoting RawPropensityWeightedVoting AdjustedPropensityWeightedVoting HighestConfidence AverageRawPropensity AverageAdjustedPropensity	アンサンブル・スコアを決定するために使用する方法を指定します。この設定は、選択された対象がフラグ型フィールドである場合にのみ適用されます。
set_ensemble_method	Voting ConfidenceWeightedVoting HighestConfidence	アンサンブル・スコアを決定するために使用する方法を指定します。この設定は、選択された対象が名義型フィールドである場合にのみ適用されます。
flag_voting_tie_selection	Random HighestConfidence RawPropensity AdjustedPropensity	票決方法が選択された場合、可否同数の解決方法を指定します。この設定は、選択された対象がフラグ型フィールドである場合にのみ適用されます。
set_voting_tie_selection	Random HighestConfidence	票決方法が選択された場合、可否同数の解決方法を指定します。この設定は、選択された対象が名義型フィールドである場合にのみ適用されます。
calculate_standard_error	フラグ	対象フィールドが連続型の場合、標準誤差の計算がデフォルトで実施され、測定された値または推定された値と真の値との差異を計算し、それらの推定がどれほど近いかを示します。

fillernode プロパティ



置換ノードで、フィールド値の置換やストレージの変更を行います。
@BLANK(@FIELD) のような、CLEM 条件に基づいて値を置換することができます。
また、すべての空白値やヌル値を特定の値に置換することもできます。置換ノードは、欠損値を置き換えるために、多くの場合データ型ノードと一緒に使用されます。

例

```
node = stream.create("filler", "My node")
node.setPropertyValue("fields", ["Age"])
node.setPropertyValue("replace_mode", "Always")
node.setPropertyValue("condition", "(\"Age\" > 60) and (\"Sex\" = \"M\")")
node.setPropertyValue("replace_with", "\"old man\"")
```

表 84. fillernode プロパティ

fillernode プロパティ	データ・タイプ	プロパティの説明
fields	リスト	検査されて置換される値のデータセットのフィールド群。
replace_mode	Always Conditional Blank Null BlankAndNull	すべての値、空白値、またはヌル値を置換できます。または、指定した条件に基づいて、置換できます。
condition	string	
replace_with	string	

filternode プロパティ



フィルターノードは、フィールドのフィルター操作（破棄）、フィールド名の変更、またはソースノードから別のノードへのフィールドのマッピングを行います。

例:

```
node = stream.create("filter", "My node")
node.setPropertyValue("default_include", True)
node.setKeyedPropertyValue("new_name", "Drug", "Chemical")
node.setKeyedPropertyValue("include", "Drug", False)
```

default_include プロパティの使用: default_include プロパティの値を設定しても、すべてのフィールドが自動的に取り込まれたり除外されたりするわけではありません。単に、現在選択されている項目に対するデフォルトが決定されるだけです。これは、「フィルター・ノード」ダイアログ・ボックスで「デフォルトでフィールドを含める」をクリックすることと、機能的に同じです。例えば、次のスクリプトを実行すると想定します。

```
node = modeler.script.stream().create("filter", "Filter")
node.setPropertyValue("default_include", False)
# Include these two fields in the list
for f in ["Age", "Sex"]:
    node.setKeyedPropertyValue("include", f, True)
```

これにより、Age (年齢) フィールドと Sex (性別) フィールドがノードを通過し、その他はすべて除外されます。前のスクリプトを実行した後に、今度は以下の行をスクリプトに追加して、さらに2つのフィールドに名前を付けるものとします。

```
node.setPropertyValue("default_include", False)
# Include these two fields in the list
for f in ["BP", "Na"]:
    node.setKeyedPropertyValue("include", f, True)
```

これにより、合計4つのフィールド (Age、Sex、BP、Na) が渡されるように、さらに2つのフィールドがフィルターに追加されます。つまり、default_include の値を False にリセットしても、すべてのフィールドが自動的にリセットされるわけではありません。

その代わり、スクリプトを使用するか「フィルター・ノード」ダイアログ・ボックス内で default_include を True にこの時点で変更すると、動作が反対になり、上記の4フィールドは上記の4フィールドは除外されます。「フィルター・ノード」ダイアログ・ボックス内のコントロールで実験することが、この相互関係を理解するうえで役に立ちます。

表 85. filternode プロパティ		
filternode プロパティ	データ・タイプ	プロパティの説明
default_include	フラグ	デフォルトの処理としてフィールドを通過させるかフィルターをかけるかの指定をするキー・プロパティ。 このプロパティを設定しても、すべてのフィールドが自動的に取り込まれたり除外されたりするわけではありません。選択したフィールドが、デフォルトでは取り込まれるか除外されるかを決めるだけです。詳細は、下の例を参照してください。
include	フラグ	フィールドを取り込むか除外するかのキー・プロパティ。
new_name	string	

historynode プロパティ



時系列ノードにより、以前レコードのフィールドのデータを含む、新規フィールドが作成されます。時系列ノードは、多くの場合、時系列データなどの継続的なデータに使用されます。履歴ノードを使用する前に、ソート・ノードを使用してデータをソートすることができます。

例

```
node = stream.create("history", "My node")
node.setPropertyValue("fields", ["Drug"])
node.setPropertyValue("offset", 1)
node.setPropertyValue("span", 3)
node.setPropertyValue("unavailable", "Discard")
node.setPropertyValue("fill_with", "undef")
```

表 86. historynode プロパティ		
historynode プロパティ	データ・タイプ	プロパティの説明
fields	リスト	履歴の対象となるフィールド。

表 86. historynode プロパティ (続き)

historynode プロパティ	データ・タイプ	プロパティの説明
offset	number	時系列フィールド値を抽出する最新レコードが、現在のレコードのいくつ前にあるかを指定します。
span	number	値を抽出する元になるレコードの前にあるレコード数を指定します。
unavailable	Discard Leave Fill	時系列として使用する前のレコードがないデータセットの先頭の数レコードを通常は指しますが、その時系列値がないレコードの取り扱い方法を指定します。
fill_with	String Number	時系列値が利用できないレコードを充填するのに使用する値 (Number) または文字列 (String) を指定します。

partitionnode プロパティ



データ区分ノードで、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に分割するデータ区分フィールドが生成されます。

例

```
node = stream.create("partition", "My node")
node.setPropertyValue("create_validation", True)
node.setPropertyValue("training_size", 33)
node.setPropertyValue("testing_size", 33)
node.setPropertyValue("validation_size", 33)
node.setPropertyValue("set_random_seed", True)
node.setPropertyValue("random_seed", 123)
node.setPropertyValue("value_mode", "System")
```

表 87. partitionnode プロパティ

partitionnode プロパティ	データ型	プロパティの説明
new_name	string	ノードにより生成されたデータ区分フィールドの名前です。
create_validation	フラグ	検証用のデータ区分を作成するかどうかを指定します。
training_size	integer	学習用区分に割り当てるレコード数のパーセンテージ (0-100)。
testing_size	integer	テスト用区分に割り当てるレコード数のパーセンテージ (0-100)。
validation_size	integer	検証用区分に割り当てるレコード数のパーセンテージ (0-100)。検証用データ区分を生成しない場合は無視されます。
training_label	string	学習用データ区分のラベル。

表 87. partitionnode プロパティ (続き)

partitionnode プロパティ	データ型	プロパティの説明
testing_label	string	テスト用データ区分のラベル。
validation_label	string	検証用データ区分のラベル。検証用データ区分を生成しない場合は無視されます。
value_mode	System SystemAndLabel Label	データ中の各データ区分を表すために使用される値を指定します。例えば、学習用サンプルは、システム整数 1、ラベル Training、またはこの 2 つを組み合わせた 1_Training のように表されます。
set_random_seed	Boolean	ユーザー指定のランダム・シードを使用するかどうかを指定します。
random_seed	integer	ユーザー定義のランダム・シードの値。この値が使用されるようにするには、set_random_seed を True に設定する必要があります。
enable_sql_generation	Boolean	SQL プッシュバックを使用してレコードをデータ区分に割り当てるかどうかを指定します。
unique_field		レコードが無作為で繰り返し可能な方法でデータ区分に割り当てるよう、入力フィールドを指定します。この値が使用されるようにするには、enable_sql_generation を True に設定する必要があります。

reclassifynode プロパティ



データ分類ノードにより、あるカテゴリー値のセットが別のセットに変換されます。データ分類ノードは、カテゴリーを縮小したり、分析対象データを再グループ化したりする場合に役立ちます。

例

```
node = stream.create("reclassify", "My node")
node.setPropertyValue("mode", "Multiple")
node.setPropertyValue("replace_field", True)
node.setPropertyValue("field", "Drug")
node.setPropertyValue("new_name", "Chemical")
node.setPropertyValue("fields", ["Drug", "BP"])
node.setPropertyValue("name_extension", "reclassified")
node.setPropertyValue("add_as", "Prefix")
node.setKeyedPropertyValue("reclassify", "drugA", True)
node.setPropertyValue("use_default", True)
node.setPropertyValue("default", "BrandX")
node.setPropertyValue("pick_list", ["BrandX", "Placebo", "Generic"])
```

表 88. reclassifynode プロパティ

reclassifynode プロパティ	データ・タイプ	プロパティの説明
mode	Single Multiple	1つのフィールドのカテゴリを再分類する場合、Singleを使用します。 Multiple (複数)を使用すると、一度に複数のフィールドを同時に変換できます。
replace_field	フラグ	
field	string	Single モードでしか使用できません。
new_name	string	Single モードでしか使用できません。
fields	[field1 field2 ... fieldn]	Multiple モードでしか使用できません。
name_extension	string	Multiple モードでしか使用できません。
add_as	Suffix Prefix	Multiple モードでしか使用できません。
reclassify	string	フィールド値用構造化プロパティ。
use_default	フラグ	デフォルト値を使用します。
default	string	デフォルト値を指定します。
pick_list	[string string ... string]	ユーザーが、既知の新しい値をインポートしてテーブル内のドロップダウン・リストをデータで埋めることができるようにします。

reordernode プロパティ



フィールド順序ノードで、下流のフィールド表示に使用する順序を定義します。この順序は、テーブル、リスト、およびフィールド・ピッカーなど、さまざまな場所のフィールドの表示に適用されます。この操作は、関心のあるフィールドを見やすくするために、幅広いデータ・セットを処理する場合に役立ちます。

例

```
node = stream.create("reorder", "My node")
node.setPropertyValue("mode", "Custom")
node.setPropertyValue("sort_by", "Storage")
node.setPropertyValue("ascending", False)
node.setPropertyValue("start_fields", ["Age", "Cholesterol"])
node.setPropertyValue("end_fields", ["Drug"])
```

表 89. reordernode プロパティ

reordernode プロパティ	データ・タイプ	プロパティの説明
mode	Custom Auto	値を自動的に並び替えたり、ユーザー指定の順序を指定することができます。

表 89. reordernode プロパティ (続き)

reordernode プロパティ	データ・タイプ	プロパティの説明
sort_by	Name Type Storage	
ascending	フラグ	
start_fields	[field1 field2 ... fieldn]	新規フィールドは、これらのフィールドの後に挿入されます。
end_fields	[field1 field2 ... fieldn]	新規フィールドは、これらのフィールドの前に挿入されます。

reprojectnode プロパティ



スポンサー モデラー では、式ビルダーの空間処理関数、時空間予測 (STP) ノード、マップ視覚化ノードなどの項目は、投影座標系を使用します。地理座標系を使用するインポート・データの座標系を変更するには、再投影ノードを使用してください。

表 90. reprojectnode プロパティ

reprojectnode プロパティ	データ・タイプ	プロパティの説明
reproject_fields	[field1 field2 ... fieldn]	再投影されるすべてのフィールドをリストします。
reproject_type	Streamdefault Specify	フィールドの再投影方法を選択します。
coordinate_system	string	フィールドに適用される座標系の名前。 例: set reprojectnode.coordinate_system = "WGS_1984_World_Mercator"

restructurenode プロパティ



再構成ノードで、名義型またはフラグ型フィールドを、これから別のフィールドの値で埋めることができるフィールドのグループへ変換します。例えば、値が貸方、現金、および借方の支払タイプという名前のフィールドがあるとすると、3つの新規フィールド(貸方、現金、借方)が作成されます。各フィールドには、実際に行われた支払いの値が含まれている可能性があります。

例

```
node = stream.create("restructure", "My node")
node.setKeyedPropertyValue("fields_from", "Drug", ["drugA", "drugX"])
node.setPropertyValue("include_field_name", True)
node.setPropertyValue("value_mode", "OtherFields")
node.setPropertyValue("value_fields", ["Age", "BP"])
```

表 91. *restructurenode* プロパティ

restructurenode プロパティ	データ・タイプ	プロパティの説明
fields_from	[カテゴリー カテゴリー] all	
include_field_name	フラグ	再構成されるフィールド名に元のフィールド名を使用するかどうかを示します。
value_mode	OtherFields Flags	再構成されるフィールドの値を指定するためのモードを示します。 OtherFields を指定すると、使用するフィールドを指定する必要があります (下を参照)。Flags を指定する場合、値は数値のフラグです。
value_fields	リスト	value_mode が OtherFields の場合は必須です。値のフィールドとして使用するフィールドを指定します。

rfmanalysisnode プロパティ



リーセンシ、フリクエンシ、マネタリー (RFM) 分析ノードを使用すると、どの顧客があなたから最後に購入したか (リーセンシ)、購入頻度 (フリクエンシ)、およびすべてのトランザクションで消費した金額 (マネタリー) を調査することにより、どの顧客が最良の顧客である可能性が高いかを定量的に判断できます。

例

```
node = stream.create("rfmanalysis", "My node")
node.setPropertyValue("recency", "Recency")
node.setPropertyValue("frequency", "Frequency")
node.setPropertyValue("monetary", "Monetary")
node.setPropertyValue("tied_values_method", "Next")
node.setPropertyValue("recalculate_bins", "IfNecessary")
node.setPropertyValue("recency_thresholds", [1, 500, 800, 1500, 2000, 2500])
```

表 92. *rfmanalysisnode* プロパティ

rfmanalysisnode プロパティ	データ・タイプ	プロパティの説明
recency	フィールド	リーセンシ フィールドを指定します。このフィールドは日付、タイムスタンプまたは単純な数値です。
frequency	フィールド	フリクエンシ フィールドを指定します。
monetary	フィールド	マネタリー・フィールドを指定します。
recency_bins	integer	生成されるリーセンシ ビンの数を指定します。
recency_weight	number	リーセンシ データに適用される重みを指定します。デフォルトは 100 です。
frequency_bins	integer	生成されるフリクエンシ ビンの数を指定します。

表 92. rfmanalysisnode プロパティ (続き)

rfmanalysisnode プロパティ	データ・タイプ	プロパティの説明
frequency_weight	number	フリクエンシ データに適用される重みを指定します。 デフォルト値は 10 です。
monetary_bins	integer	生成されるマネタリー・ビンの数を指定します。
monetary_weight	number	マネタリー・データに適用される重みを指定します。 デフォルトは 1 です。
tied_values_method	Next Current	可否同数の値のデータに配置されるビンを指定。
recalculate_bins	Always IfNecessary	
add_outliers	フラグ	recalculate_bins が IfNecessary に設定されている場合使用できます。 設定されると、下限のビンの下にあるレコードが下限のビンに追加され、上限のビンの上にあるレコードが上限のビンに追加されます。
binned_field	Recency Frequency Monetary	
recency_thresholds	value value	recalculate_bins が Always に設定されている場合使用できます。 リーセンシ ビンの上限および下限の閾値を指定します。 あるビンの上限の閾値が次のビンの下限の閾値として使用されます。例えば、[10 30 60] は、最初のビンに 10 および 30 の上限および下限の閾値があり、2 番目のビンには 30 および 60 の閾値があると定義します。
frequency_thresholds	value value	recalculate_bins が Always に設定されている場合使用できます。
monetary_thresholds	value value	recalculate_bins が Always に設定されている場合使用できます。

settoflagnode プロパティ



フラグ設定ノードで、1 つ以上の名義型フィールドに定義されたカテゴリー値に基づいた、複数のフラグ型フィールドが派生します。

例

```
node = stream.create("settoflag", "My node")
node.setKeyedPropertyValue("fields_from", "Drug", ["drugA", "drugX"])
node.setPropertyValue("true_value", "1")
node.setPropertyValue("false_value", "0")
node.setPropertyValue("use_extension", True)
node.setPropertyValue("extension", "Drug_Flag")
node.setPropertyValue("add_as", "Suffix")
node.setPropertyValue("aggregate", True)
node.setPropertyValue("keys", ["Cholesterol"])
```

表 93. settoflagnode プロパティ

settoflagnode プロパティ	データ・タイプ	プロパティの説明
fields_from	[カテゴリー カテゴリー] all	
true_value	string	フラグを設定するときにノードが使用する真 (true) の値を指定します。デフォルトは T です。
false_value	string	フラグを設定するときにノードが使用する偽 (false) の値を指定します。デフォルトは F です。
use_extension	フラグ	新規フラグ型フィールドの接尾辞または接頭辞として、拡張子を使用します。
extension	string	
add_as	Suffix Prefix	拡張子が接尾辞 (Suffix) または接頭辞 (Prefix) として追加されることを指定します。
aggregate	フラグ	キー・フィールドに基づいてレコードをグループ化します。真 (true) に設定されたレコードが 1 つでもあると、グループ内のすべてのフラグ型フィールドが有効になります。
keys	リスト	キー・フィールド。

statistictransformnode プロパティ



Statistics 変換ノードは、IBM スポス モデラー のデータ・ソースに対する IBM SPSS 統計 シンタックス・コマンドの選択を行います。このノードには、ライセンス交付を受けた IBM SPSS 統計 のコピーが必要です。

このノードのプロパティについては、[423 ページの『statistictransformnode プロパティ』](#)に記載されています。

timeintervalsnode プロパティ (廃止)



注: このノードは スポス モデラー のバージョン 18 で廃止され、新しい時系列ノードに置き換えられました。

時間区分ノードで、時系列データのモデル作成用に区分を指定し、必要に応じてラベルを作成します。値の間隔が均等に空けられていない場合は、レコード間に一律の間隔をとる必要に応じて、値を充填したり集計したりできます。

例

```
node = stream.create("timeintervals", "My node")
node.setPropertyValue("interval_type", "SecondsPerDay")
node.setPropertyValue("days_per_week", 4)
node.setPropertyValue("week_begins_on", "Tuesday")
node.setPropertyValue("hours_per_day", 10)
node.setPropertyValue("day_begins_hour", 7)
node.setPropertyValue("day_begins_minute", 5)
node.setPropertyValue("day_begins_second", 17)
node.setPropertyValue("mode", "Label")
node.setPropertyValue("year_start", 2005)
node.setPropertyValue("month_start", "January")
node.setPropertyValue("day_start", 4)
node.setKeyedPropertyValue("pad", "AGE", "MeanOfRecentPoints")
node.setPropertyValue("agg_mode", "Specify")
node.setPropertyValue("agg_set_default", "Last")
```


表 94. *timeintervalsnode* プロパティ

timeintervalsnode プロパティ	データ・タイプ	プロパティの説明
interval_type	None Periods CyclicPeriods Years Quarters Months DaysPerWeek DaysNonPeriodic HoursPerDay HoursNonPeriodic MinutesPerDay MinutesNonPeriodic SecondsPerDay SecondsNonPeriodic	
mode	Label Create	レコードに連続してラベルを付ける (Label) か、または指定された日付、タイムスタンプ、または時間フィールドに基づいて系列を構築するか (Create) を指定します。
field	フィールド	データから系列を構築する場合は、各レコードの日付または時刻を示すフィールドを指定します。
period_start	<i>integer</i>	期間または循環する期間の開始期間を指定します。
cycle_start	<i>integer</i>	循環する期間の開始サイクル。
year_start	<i>integer</i>	適用可能な区分タイプの、最初の区分が入る年。
quarter_start	<i>integer</i>	適用可能な区分タイプの、最初の区分が入る四半期。

表 94. *timeintervalsnode* プロパティ (続き)

timeintervalsnode プロパティ	データ・タイプ	プロパティの説明
month_start	January February March April May June July August September October November December	
day_start	<i>integer</i>	
hour_start	<i>integer</i>	
minute_start	<i>integer</i>	
second_start	<i>integer</i>	
periods_per_cycle	<i>integer</i>	循環する期間の、各サイクル内の期間数。
fiscal_year_begins	January February March April May June July August September October November December	四半期単位の区分の場合、会計年度が始まる月を指定します。
week_begins_on	Sunday Monday Tuesday Wednesday Thursday Friday Saturday Sunday	定期的な区分 (週当たりの日数、日当たりの時間数、日当たりの分数、日当たりの秒数) の、週が始まる曜日を指定します。

表 94. *timeintervalsnode* プロパティ (続き)

timeintervalsnode プロパティ	データ・タイプ	プロパティの説明
day_begins_hour	<i>integer</i>	定期的な区分 (日当たりの時間数、日当たりの分数、日当たりの秒数) の、日が始まる時間を指定します。 day_begins_minute と day_begins_second を組み合わせて 8:05:01 のように、正確な時刻を指定できます。下の使用例を参照してください。
day_begins_minute	<i>integer</i>	定期的な区分 (日当たりの時間数、日当たりの分数、日当たりの秒数) の、日が始まる時間の分を指定します (例えば 8:05 の 5)。
day_begins_second	<i>integer</i>	定期的な区分 (日当たりの時間数、日当たりの分数、日当たりの秒数) の、日が始まる時間の秒を指定します (例えば 08:05:17 の 17)。
days_per_week	<i>integer</i>	定期的な区分 (週当たりの日数、日当たりの時間数、日当たりの分数、日当たりの秒数) の、週当たりの日数を指定します。
hours_per_day	<i>integer</i>	定期的な区分 (日当たりの時間数、日当たりの分数、日当たりの秒数) の、1 日の時間数を指定します。
interval_increment	1 2 3 4 5 6 10 15 20 30	日当たりの分数と日当たりの秒数について、各レコード用増分の分数または秒数を指定します。
field_name_extension	<i>string</i>	
field_name_extension_as_prefix	フラグ	

表 94. timeintervalsnode プロパティ (続き)

timeintervalsnode プロパティ	データ・タイプ	プロパティの説明
date_format	"DDMMYY" "MMDDYY" "YYMMDD" "YYYYMMDD" "YYYYDDD" DAY MONTH "DD-MM-YY" "DD-MM-YYYY" "MM-DD-YY" "MM-DD-YYYY" "DD-MON-YY" "DD-MON-YYYY" "YYYY-MM-DD" "DD.MM.YY" "DD.MM.YYYY" "MM.DD.YYYY" "DD.MON.YY" "DD.MON.YYYY" "DD/MM/YY" "DD/MM/YYYY" "MM/DD/YY" "MM/DD/YYYY" "DD/MON/YY" "DD/MON/YYYY" MON YYYY q Q YYYY ww WK YYYY	
time_format	"HMMSS" "HMM" "MMSS" "HH:MM:SS" "HH:MM" "MM:SS" "(H)H:(M)M:(S)S" "(H)H:(M)M" "(M)M:(S)S" "HH.MM.SS" "HH.MM" "MM.SS" "(H)H.(M)M.(S)S" "(H)H.(M)M" "(M)M.(S)S"	
aggregate	Mean Sum Mode Min Max First Last TrueIfAnyTrue	フィールドの集計方法を指定します。

表 94. *timeintervalsnode* プロパティ (続き)

timeintervalsnode プロパティ	データ・タイプ	プロパティの説明
pad	Blank MeanOfRecentPoints True False	フィールドの充填方法を指定します。
agg_mode	All Specify	必要に応じてデフォルトの関数ですべてのフィールドを集計または充填するかどうかを指定します。または、使用するフィールドと関数を指定します。
agg_range_default	Mean Sum Mode Min Max	連続型フィールドを集計するときに使用するデフォルトの関数を指定します。
agg_set_default	Mode First Last	名義型フィールドを集計するときに使用するデフォルトの関数を指定します。
agg_flag_default	TrueIfAnyTrue Mode First Last	
pad_range_default	Blank MeanOfRecentPoints	連続型フィールドをパディングするときに使用するデフォルトの関数を指定します。
pad_set_default	Blank MostRecentValue	
pad_flag_default	Blank True False	

表 94. *timeintervalsnode* プロパティ (続き)

timeintervalsnode プロパティ	データ・タイプ	プロパティの説明
max_records_to_create	<i>integer</i>	系列を充填するときに作成する最大レコード数を指定します。
estimation_from_beginning	フラグ	
estimation_to_end	フラグ	
estimation_start_offset	<i>integer</i>	
estimation_num_holdouts	<i>integer</i>	
create_future_records	フラグ	
num_future_records	<i>integer</i>	
create_future_field	フラグ	
future_field_name	<i>string</i>	

transposenode プロパティ



行列入替ノードは、レコードがフィールドになり、フィールドがレコードになるように、行内と列内のデータを交換します。

例

```
node = stream.create("transpose", "My node")
node.setPropertyValue("transposed_names", "Read")
node.setPropertyValue("read_from_field", "TimeLabel")
node.setPropertyValue("max_num_fields", "1000")
node.setPropertyValue("id_field_name", "ID")
```

表 95. *transposenode* プロパティ

transposenode プロパティ	データ・タイプ	プロパティの説明
transpose_method	<i>enum</i>	行列入替方法を、「通常」(normal)、「CASE から VAR へ」(casetovar)、「VAR から CASE へ」(vartocase)の中から指定します。
transposed_names	Prefix Read	「通常」行列入替方法のプロパティ。新しいフィールド名は、指定された接頭辞 (Prefix) に基づいて自動的に作成できます。または、既存のデータ内のフィールドからフィールド名を読み込むことができます (Read)。
prefix	<i>string</i>	「通常」行列入替方法のプロパティ。
num_new_fields	<i>integer</i>	「通常」行列入替方法のプロパティ。接頭辞を使用する場合は、作成する新しいフィールドの最大数を指定します。

表 95. *transposenode* プロパティ (続き)

transposenode プロパティ	データ・タイプ	プロパティの説明
<code>read_from_field</code>	フィールド	「通常」行列入替方法のプロパティ。名前が読み込まれるフィールド。これはインスタンス化されたフィールドであることが必要です。そうでない場合は、ノードが実行されるとときにエラーが発生します。
<code>max_num_fields</code>	<i>integer</i>	「通常」行列入替方法のプロパティ。フィールドから名前を読み込む場合は、異常に大量のフィールドを作成しないように、フィールド数の上限を指定します。
<code>transpose_type</code>	Numeric String Custom	「通常」行列入替方法のプロパティ。デフォルトでは連続型のフィールドのみの行列が入れ替えられますが、代わりに、数値フィールドのカスタム (ユーザー設定) サブセットを選択またはすべての文字列フィールドを入れ替えることもできます。
<code>transpose_fields</code>	リスト	「通常」行列入替方法のプロパティ。Custom (ユーザー設定) オプションを使用するときに、行列を入れ替えるフィールドを指定します。
<code>id_field_name</code>	フィールド	「通常」行列入替方法のプロパティ。
<code>transpose_casetovar_idfields</code>	フィールド	「CASE から VAR へ (casetovar)」行列入替方法のプロパティ。複数のフィールドをインデックス・フィールドとして使用することを受け入れます。 field1 ... fieldN
<code>transpose_casetovar_columnfields</code>	フィールド	「CASE から VAR へ (casetovar)」行列入替方法のプロパティ。複数のフィールドを列フィールドとして使用することを受け入れます。 field1 ... fieldN
<code>transpose_casetovar_valuefields</code>	フィールド	「CASE から VAR へ (casetovar)」行列入替方法のプロパティ。複数のフィールドを値フィールドとして使用することを受け入れます。 field1 ... fieldN
<code>transpose_vartocase_idfields</code>	フィールド	「VAR から CASE へ (vartocase)」行列入替方法のプロパティ。複数のフィールドを ID 変数フィールドとして使用することを受け入れます。 field1 ... fieldN

表 95. *transposenode* プロパティ (続き)

transposenode プロパティ	データ・タイプ	プロパティの説明
transpose_vartocase_valfields	フィールド	「VAR から CASE へ (vartocase)」行列入替方法のプロパティ。複数のフィールドを値変数フィールドとして使用することを受け入れます。 field1 ... fieldN

typenode プロパティ



データ型ノードで、フィールドのメタデータとプロパティを指定します。例えば、各フィールドに、測定の尺度 (連続型、名義型、順序型、またはフラグ) を指定し、欠損値とシステムヌルの処理のためのオプションを設定し、モデル作成の目的に対するフィールドの役割を設定し、フィールドと値のラベルを指定し、フィールドの値を指定します。

例

```
node = stream.createAt("type", "My node", 50, 50)
node.setKeyedPropertyValue("check", "Cholesterol", "Coerce")
node.setKeyedPropertyValue("direction", "Drug", "Input")
node.setKeyedPropertyValue("type", "K", "Range")
node.setKeyedPropertyValue("values", "Drug", ["drugA", "drugB", "drugC",
"drugD", "drugX",
"drugY", "drugZ"])
node.setKeyedPropertyValue("null_missing", "BP", False)
node.setKeyedPropertyValue("whitespace_missing", "BP", False)
node.setKeyedPropertyValue("description", "BP", "Blood Pressure")
node.setKeyedPropertyValue("value_labels", "BP", [["HIGH", "High Blood
Pressure"],
["NORMAL", "normal blood pressure"]])
```

ある種の場合、ほかのノードが正しく機能するように、フラグ設定ノードの **fields from** プロパティのように、データ型ノードを完全にインスタンス化する必要がある場合があります。フィールドをインスタンス化するには、次のように、テーブル・ノードを接続して実行するだけです。

```
tablenode = stream.createAt("table", "Table node", 150, 50)
stream.link(node, tablenode)
tablenode.run(None)
stream.delete(tablenode)
```


表 96. <i>typenode</i> プロパティ		
typenode プロパティ	データ・タイプ	プロパティの説明
direction	Input Target Both None Partition Split Frequency RecordID	フィールドの役割のキープロパティ。 注: 値 In と Out は廃止されました。今後のリリースではサポートが中断される場合があります。
type	Range Flag Set Typeless Discrete OrderedSet Default	フィールドの測定の尺度 (以前はフィールドの「タイプ」と呼ばれます)。 type を Default に設定すると、すべての values パラメーター設定をクリアされ、および value_mode の値が Specify の場合は、Read にリセットされます。 value_mode が Pass または Read に設定されている場合、type の設定は、value_mode には影響しません。 注: 内部で使用するデータ型は、データ型ノードに表示されるデータ型とは異なります。次のように対応します: 範囲型 > 連続セット型 -> 名義順序セット型 -> 順序離散型 -> カテゴリー型
storage	Unknown String Integer Real Time Date Timestamp	フィールドのストレージ・タイプ用読み込み専用キー・プロパティ。

表 96. *typenode* プロパティ (続き)

typenode プロパティ	データ・タイプ	プロパティの説明
check	None Nullify Coerce Discard Warn Abort	フィールド・タイプと範囲の検査用のキー・プロパティ。
values	[値 値]	連続型フィールドの場合、最初の値が最小値で最後の値が最大値になります。名義型フィールドの場合、すべての値を指定します。フラグ型の場合、最初の値が <i>false</i> (偽) を、最後の値が <i>true</i> (真) を表します。このプロパティを設定すると、 <i>value_mode</i> プロパティの値が自動的に <i>Specify</i> に設定されます。
value_mode	Read Pass Read+ Current Specify	値の設定方法を決定します。このプロパティに <i>Specify</i> を直接には設定できないことに注意してください。特定の値を使用するには、 <i>values</i> プロパティを設定します。
extend_values	フラグ	<i>value_mode</i> が <i>Read</i> に設定された場合に適用されます。新しく読み込んだ値を、フィールドの既存の値に追加する場合は、 <i>T</i> を設定します。新しく読み込んだ値を優先して、既存の値を破棄する場合は、 <i>F</i> を設定します。
enable_missing	フラグ	<i>T</i> を設定した場合、フィールドの欠損値の追跡が有効になります。
missing_values	[value value ...]	欠損データを示すデータ値を指定します。
range_missing	フラグ	フィールドに欠損値 (空白) の範囲が定義されているかどうかを指定します。
missing_lower	string	<i>range_missing</i> が真 (<i>true</i>) の場合、欠損値範囲の下限値を指定します。
missing_upper	string	<i>range_missing</i> が真 (<i>true</i>) の場合、欠損値範囲の上限値を指定します。

表 96. <i>typenode</i> プロパティ (続き)		
typenode プロパティ	データ・タイプ	プロパティの説明
null_missing	フラグ	Tを設定した場合、ヌル値 (ソフトウェアでは \$null\$ として表示される未定義値) は欠損値と見なされます。
whitespace_missing	フラグ	Tを設定した場合、空白類 (スペース、タブ、および改行) だけを含む値が欠損値と見なされます。
description	string	フィールドの説明を指定します。
value_labels	[[Value LabelString] [Value LabelString] ...]	値のペアのためのラベルを指定します。
display_places	integer	フィールドが表示されるとき的小数部の桁数を設定します (REAL ストレージのフィールドにのみ適用)。値-1 は、ストリームのデフォルトを使用します。
export_places	integer	フィールドが出力されるとき的小数部の桁数を設定します (REAL ストレージのフィールドにのみ適用)。値-1 は、ストリームのデフォルトを使用します。
decimal_separator	DEFAULT PERIOD COMMA	フィールドの小数点記号を指定します (REAL ストレージのフィールドにのみ適用)。
date_format	"DDMMYY" "MMDDYY" "YYMMDD" "YYYYMMDD" "YYYYDDD" DAY MONTH "DD-MM-YY" "DD-MM-YYYY" "MM-DD-YY" "MM-DD-YYYY" "DD-MON-YY" "DD-MON-YYYY" "YYYY-MM-DD" "DD.MM.YY" "DD.MM.YYYY" "MM.DD.YYYY" "DD.MON.YY" "DD.MON.YYYY" "DD/MM/YY" "DD/MM/YYYY" "MM/DD/YY" "MM/DD/YYYY" "DD/MON/YY" "DD/MON/YYYY" MON YYYY q Q YYYY ww WK YYYY	フィールドの日付形式を設定します (DATE または TIMESTAMP ストレージのフィールドにのみ適用されます)。

表 96. *typenode* プロパティ (続き)

typenode プロパティ	データ・タイプ	プロパティの説明
time_format	"HHMMSS" "HHMM" "MMSS" "HH:MM:SS" "HH:MM" "MM:SS" "(H)H:(M)M:(S)S" "(H)H:(M)M" "(M)M:(S)S" "HH.MM.SS" "HH.MM" "MM.SS" "(H)H.(M)M.(S)S" "(H)H.(M)M" "(M)M.(S)S"	フィールドの日付形式を設定します (TIME または TIMESTAMP ストレージのフィールドにのみ適用されます)。
number_format	DEFAULT STANDARD SCIENTIFIC CURRENCY	フィールドに数値の表示形式を設定します。
standard_places	<i>integer</i>	フィールドが標準形式で表示されるときの小数点以下の桁数を指定します。値-1は、ストリームのデフォルトを使用します。既存の <i>display_places</i> スロットでもこの設定が変更されますが、現在は廃止されています。
scientific_places	<i>integer</i>	フィールドが科学系の形式で表示されるときの小数点以下の桁数を設定します。値-1は、ストリームのデフォルトを使用します。
currency_places	<i>integer</i>	フィールドが通貨の形式で表示されるときフィールドの小数点以下の桁数を設定します。値-1は、ストリームのデフォルトを使用します。
grouping_symbol	DEFAULT NONE LOCALE PERIOD COMMA SPACE	フィールドにグループ化シンボルを設定します。
column_width	<i>integer</i>	フィールドに列幅を設定します。値-1は、列の幅が Auto に設定されます。

表 96. *typenode* プロパティ (続き)

typenode プロパティ	データ・タイプ	プロパティの説明
justify	AUTO CENTER LEFT RIGHT	フィールドに列調整を設定します。
measure_type	Range / MeasureType.RANGE Discrete / MeasureType.DISCRETE Flag / MeasureType.FLAG Set / MeasureType.SET OrderedSet / MeasureType.ORDERED_SET Typeless / MeasureType.TYPELESS Collection / MeasureType.COLLECTION Geospatial / MeasureType.GEOSPATIAL	このキー付きプロパティは、フィールドに関連付けられた尺度を定義するために使用できるという点で、 <code>type</code> と類似しています。異なるのは、Python スクリプトで、 <code>getter</code> 関数が常に <code>MeasureType</code> 値を返す一方で、 <code>setter</code> 関数に <code>MeasureType</code> 値のうちの 1 つを渡すこともできるという点です。
collection_measure	Range / MeasureType.RANGE Flag / MeasureType.FLAG Set / MeasureType.SET OrderedSet / MeasureType.ORDERED_SET Typeless / MeasureType.TYPELESS	収集フィールド (深さが 0 のリスト) の場合、このキー付きプロパティは、基礎となる値に関連付けられた尺度タイプを定義します。

表 96. *typenode* プロパティ (続き)

typenode プロパティ	データ・タイプ	プロパティの説明
geo_type	Point MultiPoint LineString MultiLineString Polygon MultiPolygon	地理空間フィールドの場合、このキー付きプロパティにより、このフィールドが表す地理空間オブジェクトのタイプが定義されます。これは、値のリストの深さと整合している必要があります。
has_coordinate_system	<i>Boolean</i>	地理空間フィールドの場合、このプロパティにより、このフィールドに座標系があるかどうか定義されます。
coordinate_system	<i>string</i>	地理空間フィールドの場合、このキー付きプロパティにより、このフィールドの座標系が定義されます。
custom_storage_type	Unknown / MeasureType.UNKNOWN String / MeasureType.STRING Integer / MeasureType.INTEGER Real / MeasureType.REAL Time / MeasureType.TIME Date / MeasureType.DATE Timestamp / MeasureType.TIMESTAMP List / MeasureType.LIST	このキー付きプロパティは、フィールドのオーバーライドストレージを定義するために使用できるという点で、 <i>custom_storage</i> と類似しています。異なるのは、Python スクリプトで、 <i>getter</i> 関数が常に <i>StorageType</i> 値を返す一方で、 <i>setter</i> 関数に <i>StorageType</i> 値のうちの 1 つを渡すこともできるという点です。
custom_list_storage_type	String / MeasureType.STRING Integer / MeasureType.INTEGER Real / MeasureType.REAL Time / MeasureType.TIME Date / MeasureType.DATE Timestamp / MeasureType.TIMESTAMP	リストフィールドの場合、このキー付きプロパティにより、基礎となる値のストレージタイプが指定されます。

表 96. <i>typenode</i> プロパティ (続き)		
typenode プロパティ	データ・タイプ	プロパティの説明
<code>custom_list_depth</code>	<i>integer</i>	リスト フィールドの場合、このキー付きプロパティにより、フィールドの深さが指定されます。
<code>max_list_length</code>	<i>integer</i>	地理空間 または集合 のいずれかの尺度を持つデータのみで使用できます。 リストの最大長を設定するには、リストに入れることができる要素の数を指定します。
<code>max_string_length</code>	<i>integer</i>	データ型不明 のデータでのみ使用可能で、SQL を生成してテーブルを作成するときに使用されます。 データの最大文字列の値を入力します。これにより、テーブルに生成される列が、その文字列を含めるのに十分な大きさになります。

第 12 章 グラフ作成ノードのプロパティ

グラフ作成ノードの共通プロパティ

このセクションでは、グラフ作成ノードで利用できるプロパティについて、共通なプロパティとノード・タイプ固有のプロパティも含めて説明します。

表 97. グラフ作成ノードの共通プロパティ		
グラフ作成ノードの共通プロパティ	データ・タイプ	プロパティの説明
title	string	タイトルを指定します。 例:"This is a title."
caption	string	解説を指定します。 例:"This is a caption."
output_mode	Screen File	グラフ作成ノードからの出力が表示されるか、ファイルへ書き込まれるかを指定します。
output_format	BMP JPEG PNG HTML output (.cou)	出力のタイプを指定します。 出力可能なタイプは、各ノードに応じて変化します。
full_filename	string	グラフ作成ノードから生成されたグラフの、出力先のパスとファイル名を指定します。
use_graph_size	フラグ	下に説明する幅と高さのプロパティを使用してグラフのサイズが明示して設定されるかどうかを制御します。 画面に出力されるグラフにだけ影響します。 棒グラフ・ノードには使用できません。
graph_width	number	use_graph_size が True の場合、グラフの幅をピクセル数で指定します。
graph_height	number	use_graph_size が True の場合、グラフの高さをピクセル数で指定します。

オプション フィールドの無効化

散布図のオーバーレイ・フィールドなどのオプション・フィールドは、次の例のようにプロパティ値に " " (空文字列) を設定することにより、無効化することができます。

```
plotnode.setPropertyValue("color_field", "")
```

色の指定

表題、解説、背景、およびラベルの色は、ハッシュ記号(#)で始まる 16 進文字列で指定することができます。例えば、グラフの背景を空色にするには、次の文を指定します。

```
mygraphnode.setPropertyValue("graph_background", "#87CEEB")
```

ここで、最初の 2 桁 87 は赤色の量を、次の 2 桁 CE は緑の量を、最後の 2 桁 EB は青の量を示します。各桁は、0 から 9 または A から F の範囲の値になります。これらの値を使用して、赤-緑-青 (RGB) の色を指定します。

注: 色を RGB で指定する場合、ユーザー インターフェースのフィールド ピッカーを使用して正しい色コードを決定することができます。ピッカーを目的の色の上にかざせば、その色コードがツールヒントに表示されます。

collectionnode プロパティ



コレクション・ノードには、ある数値フィールドの値の、別の数値フィールドの値に対する分布が表示されます。(ヒストグラムに似たグラフを作成します。) 集計棒グラフは、値が時間の経過とともに変化する変数やフィールドを表示する場合に役立ちます。3 次元グラフを使用して、分布をカテゴリー別に表示するシンボル値軸を追加することもできます。

例

```
node = stream.create("collection", "My node")
# "Plot" tab
node.setPropertyValue("three_D", True)
node.setPropertyValue("collect_field", "Drug")
node.setPropertyValue("over_field", "Age")
node.setPropertyValue("by_field", "BP")
node.setPropertyValue("operation", "Sum")
# "Overlay" section
node.setPropertyValue("color_field", "Drug")
node.setPropertyValue("panel_field", "Sex")
node.setPropertyValue("animation_field", "")
# "Options" tab
node.setPropertyValue("range_mode", "Automatic")
node.setPropertyValue("range_min", 1)
node.setPropertyValue("range_max", 100)
node.setPropertyValue("bins", "ByNumber")
node.setPropertyValue("num_bins", 10)
node.setPropertyValue("bin_width", 5)
```

表 98. collectionnode プロパティ

collectionnode プロパティ	データ・タイプ	プロパティの説明
over_field	フィールド	
over_label_auto	フラグ	
over_label	string	
collect_field	フィールド	
collect_label_auto	フラグ	
collect_label	string	
three_D	フラグ	
by_field	フィールド	

表 98. collectionnode プロパティ (続き)

collectionnode プロパティ	データ・タイプ	プロパティの説明
by_label_auto	フラグ	
by_label	string	
operation	Sum Mean Min Max SDev	
color_field	string	
panel_field	string	
animation_field	string	
range_mode	Automatic UserDefined	
range_min	number	
range_max	number	
bins	ByNumber ByWidth	
num_bins	number	
bin_width	number	
use_grid	フラグ	
graph_background	COLOR	標準のグラフ色は、このセクションの最初に説明されています。
page_background	COLOR	標準のグラフ色は、このセクションの最初に説明されています。

distributionnode プロパティ



棒グラフ・ノードで、ローンの種類や性別など、シンボル値 (カテゴリー) の出現頻度を表示します。通常、棒グラフ・ノードを使用してデータの不均衡を表示しますが、そのデータはモデルの作成前にバランス・ノードを使用して修正できます。

例

```
node = stream.create("distribution", "My node")
# "Plot" tab
node.setPropertyValue("plot", "Flags")
node.setPropertyValue("x_field", "Age")
node.setPropertyValue("color_field", "Drug")
```

```
node.setPropertyValue("normalize", True)
node.setPropertyValue("sort_mode", "ByOccurence")
node.setPropertyValue("use_proportional_scale", True)
```

表 99. *distributionnode* プロパティ

distributionnode プロパティ	データ・タイプ	プロパティの説明
plot	SelectedFields Flags	
x_field	フィールド	
color_field	フィールド	オーバーレイ・フィールド。
normalize	フラグ	
sort_mode	ByOccurence Alphabetic	
use_proportional_scale	フラグ	

evaluationnode プロパティ



評価ノードは、予測モデルの評価と比較に用いられます。評価グラフで、モデルが特定の結果をどの程度予測するかを表示します。それによって、予測値と予測の確信度に基づいたレコードがソートされます。そして、レコードが等サイズ(分位)のグループに分割され、各分位のビジネスに関する基準の値が、高い方から降順で作図されます。プロットには、複数のモデルが異なる線で示されます。

例

```
node = stream.create("evaluation", "My node")
# "Plot" tab
node.setPropertyValue("chart_type", "Gains")
node.setPropertyValue("cumulative", False)
node.setPropertyValue("field_detection_method", "Name")
node.setPropertyValue("inc_baseline", True)
node.setPropertyValue("n_tile", "Deciles")
node.setPropertyValue("style", "Point")
node.setPropertyValue("point_type", "Dot")
node.setPropertyValue("use_fixed_cost", True)
node.setPropertyValue("cost_value", 5.0)
node.setPropertyValue("cost_field", "Na")
node.setPropertyValue("use_fixed_revenue", True)
node.setPropertyValue("revenue_value", 30.0)
node.setPropertyValue("revenue_field", "Age")
node.setPropertyValue("use_fixed_weight", True)
node.setPropertyValue("weight_value", 2.0)
node.setPropertyValue("weight_field", "K")
```

表 100. evaluationnode プロパティ

evaluationnode プロパティ	データ・タイプ	プロパティの説明
chart_type	Gains Response Lift Profit ROI ROC	
inc_baseline	フラグ	
field_detection_method	Metadata Name	
use_fixed_cost	フラグ	
cost_value	number	
cost_field	string	
use_fixed_revenue	フラグ	
revenue_value	number	
revenue_field	string	
use_fixed_weight	フラグ	
weight_value	number	
weight_field	フィールド	
n_tile	Quartiles Quintiles Deciles Vingtiles Percentiles 1000-tiles	
cumulative	フラグ	
style	Line Point	

表 100. evaluationnode プロパティ (続き)

evaluationnode プロパティ	データ・タイプ	プロパティの説明
point_type	Rectangle Dot Triangle Hexagon Plus Pentagon Star BowTie HorizontalDash VerticalDash IronCross Factory House Cathedral OnionDome ConcaveTriangle OblateGlobe CatEye FourSidedPillow RoundRectangle Fan	
export_data	フラグ	
data_filename	string	
delimiter	string	
new_line	フラグ	
inc_field_names	フラグ	
inc_best_line	フラグ	
inc_business_rule	フラグ	
business_rule_condition	string	
plot_score_fields	フラグ	
score_fields	[field1 ... fieldN]	
target_field	フィールド	
use_hit_condition	フラグ	
hit_condition	string	
use_score_expression	フラグ	
score_expression	string	
caption_auto	フラグ	

graphboardnode プロパティ



グラフボード・ノードでは、単一のノードにさまざまな種類のグラフを提供しています。このノードを使用して、検証するデータ・フィールドを選択し、選択したデータに使用できるグラフを選択できます。選択したフィールドに適さないグラフの種類は、ノードによって自動的に除外されます。

注: グラフ タイプに対して無効なプロパティを設定した場合 (例えば、ヒストグラフに対して y_field を指定した場合)、そのプロパティは無視されます。

注: UI には、さまざまなグラフ タイプの「詳細」タブに「要約」フィールドがあります。このフィールドは、現在スクリプトではサポートされていません。

例

```
node = stream.create("graphboard", "My node")
node.setPropertyValue("graph_type", "Line")
node.setPropertyValue("x_field", "K")
node.setPropertyValue("y_field", "Na")
```

表 101. graphboardnode プロパティ		
graphboard プロパティ	データ・タイプ	プロパティの説明
graph_type	2DDotplot 3DArea 3DBar 3DDensity 3DHistogram 3DPie 3DScatterplot Area ArrowMap Bar BarCounts BarCountsMap BarMap BinnedScatter Boxplot Bubble ChoroplethMeans ChoroplethMedians ChoroplethSums ChoroplethValues	グラフの種類を識別します。

表 101. *graphboardnode* プロパティ (続き)

graphboard プロパティ	データ・タイプ	プロパティの説明
	ChoroplethCounts	
	CoordinateMap	
	CoordinateChoroplethMeans	
	CoordinateChoroplethMedians	
	CoordinateChoroplethSums	
	CoordinateChoroplethValues	
	CoordinateChoroplethCounts	
	Dotplot	
	Heatmap	
	HexBinScatter	
	Histogram	
	Line	
	LineChartMap	
	LineOverlayMap	
	Parallel	
	Path	
	Pie	
	PieCountMap	
	PieCounts	
	PieMap	

表 101. graphboardnode プロパティ (続き)

graphboard プロパティ	データ・タイプ	プロパティの説明
	PointOverlayMap PolygonOverlayMap Ribbon Scatterplot SPLOM Surface	
x_field	フィールド	x 軸のカスタム (ユーザー設定) ラベルを指定します。ラベルでのみ使用できます。
y_field	フィールド	y 軸のカスタム (ユーザー設定) ラベルを指定します。ラベルでのみ使用できます。
z_field	フィールド	3 次元グラフの一部で使します。
color_field	フィールド	ヒート・マップで使します。
size_field	フィールド	バブル・プロットで使します。
categories_field	フィールド	
values_field	フィールド	
rows_field	フィールド	
columns_field	フィールド	
fields	フィールド	
start_longitude_field	フィールド	参照マップの矢印で使します。
end_longitude_field	フィールド	
start_latitude_field	フィールド	
end_latitude_field	フィールド	
data_key_field	フィールド	さまざまなマップで使します。
panelrow_field	string	
panelcol_field	string	
animation_field	string	

表 101. graphboardnode プロパティ (続き)

graphboard プロパティ	データ・タイプ	プロパティの説明
longitude_field	フィールド	マップ上の座標で使します。
latitude_field	フィールド	
map_color_field	フィールド	

histogramnode プロパティ



ヒストグラム・ノードでは、数値フィールドの値の出現頻度が示されます。多くの場合、ヒストグラム・ノードは、操作やモデルの構築前にデータを調べるために使われます。棒グラフ・ノードと同様、ヒストグラム・ノードにより、データ内の不均衡がしばしば明らかになります。

例

```
node = stream.create("histogram", "My node")
# "Plot" tab
node.setPropertyValue("field", "Drug")
node.setPropertyValue("color_field", "Drug")
node.setPropertyValue("panel_field", "Sex")
node.setPropertyValue("animation_field", "")
# "Options" tab
node.setPropertyValue("range_mode", "Automatic")
node.setPropertyValue("range_min", 1.0)
node.setPropertyValue("range_max", 100.0)
node.setPropertyValue("num_bins", 10)
node.setPropertyValue("bin_width", 10)
node.setPropertyValue("normalize", True)
node.setPropertyValue("separate_bands", False)
```

表 102. histogramnode プロパティ

histogramnode プロパティ	データ・タイプ	プロパティの説明
field	フィールド	
color_field	フィールド	
panel_field	フィールド	
animation_field	フィールド	
range_mode	Automatic UserDefined	
range_min	<i>number</i>	
range_max	<i>number</i>	
bins	ByNumber ByWidth	
num_bins	<i>number</i>	
bin_width	<i>number</i>	

表 102. histogramnode プロパティ (続き)

histogramnode プロパティ	データ・タイプ	プロパティの説明
normalize	フラグ	
separate_bands	フラグ	
x_label_auto	フラグ	
x_label	string	
y_label_auto	フラグ	
y_label	string	
use_grid	フラグ	
graph_background	COLOR	標準のグラフ色は、このセクションの最初に説明されています。
page_background	COLOR	標準のグラフ色は、このセクションの最初に説明されています。
normal_curve	フラグ	正規分布のカーブを出力に表示するかどうかを指定します。

mapvisualization プロパティ



マップ視覚化ノードは、複数の入力接続を受け入れて、地理空間データを一連の層としてマップに表示することができます。各層は単一の地理空間フィールドです。例えば、基本層を国のマップとし、その上に道路の層、川の層、町の層を設けることができます。

表 103. mapvisualization プロパティ

mapvisualization プロパティ	データ・タイプ	プロパティの説明
tag	string	入力用のタグの名前を設定します。デフォルトのタグは、入力がノードに接続された順序に基づく数値です (最初の接続タグは 1、2 番目の接続タグは 2 という方法で数値が設定されます)。

表 103. mapvisualization プロパティ (続き)

mapvisualization プロパティ	データ・タイプ	プロパティの説明
layer_field	フィールド	マップ上に層として表示する、データ・セットからの地理フィールドを選択します。デフォルトの選択内容は、次のソート順に基づきます。 • 最初 - 点 • 行ストリング • 多角形 • 複数点 • 複数行ストリング • 最後 - 多角形群 同じ尺度タイプを持つ 2 つのフィールドがある場合、デフォルトでは、名前のアルファベット順で最初のフィールドが選択されます。
color_type	Boolean	標準の色を地理フィールドのすべてのフィーチャーに適用するか、オーバーレイ フィールドを適用するかを指定します (オーバーレイ フィールドでは、データ・セットの他のフィールドの値に基づいて、フィーチャーに色が付けられます)。指定できる値は、standard または overlay です。デフォルトは standard です。
color	string	color_type に standard を選択した場合、ドロップダウンには、「ユーザー・オプション」の「表示」タブにある「グラフ カテゴリーの色順序」と同じ色パレットが含まれます。 デフォルトの「グラフ カテゴリーの色」は、1 です。
color_field	フィールド	color_type に overlay を選択した場合、ドロップダウンには、層として選択された地理フィールドと同じデータ・セットからのすべてのフィールドが含まれます。
symbol_type	Boolean	標準の記号を地理フィールドのすべてのレコードに適用するか、オーバーレイ記号を適用するかを指定します (オーバーレイ記号では、データ・セットの他のフィールドの値に基づいて、ポイントの記号アイコンが変更されます)。指定できる値は、standard または overlay です。デフォルトは standard です。

表 103. mapvisualization プロパティ (続き)

mapvisualization プロパティ	データ・タイプ	プロパティの説明
symbol	string	symbol_type に standard を選択した場合、ドロップダウンには、マップ上にポイントを表示するために使用される記号の選択項目が含まれます。
symbol_field	フィールド	symbol_type に overlay を選択した場合、ドロップダウンには、層として選択された地理フィールドと同じデータ・セットからの名義型フィールド、順序型フィールド、またはカテゴリー型フィールドがすべて含まれます。
size_type	Boolean	標準のサイズを地理フィールドのすべてのレコードに適用するか、オーバーレイのサイズを適用するかを指定します (オーバーレイのサイズでは、データ・セットの他のフィールドの値に基づいて、記号アイコンのサイズまたは線の太さが変更されます)。指定できる値は、standard または overlay です。デフォルトは standard です。
size	string	size_type、point、または multipoint に standard を選択した場合、ドロップダウンには、選択した記号のサイズの選択項目が含まれます。linestring または multilinesring の場合、ドロップダウンには、線の太さの選択項目が含まれます。
size_field	フィールド	size_type に overlay を選択した場合、ドロップダウンには、層として選択された地理フィールドと同じデータ・セットからのすべてのフィールドが含まれます。
transp_type	Boolean	標準の透過度を地理フィールドのすべてのレコードに適用するか、オーバーレイ透過度を適用するかを指定します (オーバーレイ透過度では、データ・セットの他のフィールドの値に基づいて、記号、線、または多角形の透過度のレベルが変更されます)。指定できる値は、standard または overlay です。デフォルトは standard です。

表 103. mapvisualization プロパティ (続き)

mapvisualization プロパティ	データ・タイプ	プロパティの説明
transp	integer	transp_type に standard を選択した場合、ドロップダウンには、透過度レベルの選択項目が含まれます。この項目は、0% (不透明) から 100% (透明) まで 10% 刻みで増加します。マップ上のポイント、線、または多角形の透過度を設定します。 size_type に overlay を選択した場合、ドロップダウンには、層として選択された地理フィールドと同じデータ・セットからのすべてのフィールドが含まれます。 points、multipoints、linestrings、および multilinstrings、polygons および multipolygons (最下層) の場合、デフォルトは 0% です。一番下の層ではない polygons および multipolygons の場合、デフォルトは 50% です (これらの多角形の下に層が表示されないようにするため)。
transp_field	フィールド	transp_type に overlay を選択した場合、ドロップダウンには、層として選択された地理フィールドと同じデータ・セットからのすべてのフィールドが含まれます。
data_label_field	フィールド	マップのデータ ラベルとして使用するフィールドを指定します。例えば、この設定の適用先の層が多角形の層の場合は、データ ラベルを name フィールドにして、それぞれの多角形の名前を含めることができます。そのため、ここで name フィールドを選択すると、それらの名前がマップに表示されるようになります。
use_hex_binning	Boolean	六角ビン分割を有効にし、すべての集計ドロップダウンを有効にします。デフォルトでは、この設定はオフになっています。

表 103. mapvisualization プロパティ (続き)

mapvisualization プロパティ	データ・タイプ	プロパティの説明
color_aggregation および transp_aggregation	string	六角ビン分割を使用するポイント層に対して「オーバーレイ」フィールドを選択した場合は、六角形の中にあるすべてのポイントについて、そのフィールドのすべての値を集計する必要があります。したがって、マップに適用するすべての「オーバーレイ」フィールドについて、集計関数を指定する必要があります。 使用可能な集計関数を次に示します。 連続型 (実数または整数のストレージ): • 合計 • 平均値 • 最小 • 最大 • 中央値 • 第 1 四分位数 • 第 3 四分位数 連続型 (時間、日付、またはタイムスタンプのストレージ): • 平均値 • 最小 • 最大 名義型/カテゴリ型: • モード • 最小 • 最大 フラグ型: • いずれかが真の場合は真 • いずれかが偽の場合は偽
custom_storage	string	フィールド全体のストレージ タイプを設定します。デフォルトは List です。List を指定した場合は、次の custom_value_storage コントロールと list_depth コントロールが無効になります。
custom_value_storage	string	フィールド全体ではなく、リスト内の要素のストレージ タイプを設定します。デフォルトは Real です。

表 103. mapvisualization プロパティ (続き)

mapvisualization プロパティ	データ・タイプ	プロパティの説明
list_depth	integer	リスト フィールドの深さを設定します。必要な深さは、地理フィールドのタイプによって異なり、次の基準に従います。 • ポイント - 0 • 行ストリング - 1 • 多角形 - 2 • 複数点 - 1 • 複数行ストリング - 2 • 多角形群 - 3 リストに変換し直す地理空間フィールドのタイプと、その種類のフィールドに必要な深さを把握しておく必要があります。設定に誤りがあると、フィールドを使用できません。 デフォルト値は 0、最小値は 0、最大値は 10 です。

multiplotnode プロパティ



線グラフ・ノードでは、1つの X フィールドに対して複数の Y フィールドを表示する作図が作成されます。Y フィールドは色付きの線で作図され、それぞれ「スタイル」フィールドを「ライン」に、「X モード」フィールドを「ソート」に設定した散布図ノードに相当します。線グラフは、複数の変数の変動を長期にわたって調査するときに役立ちます。

例

```
node = stream.create("multiplot", "My node")
# "Plot" tab
node.setPropertyValue("x_field", "Age")
node.setPropertyValue("y_fields", ["Drug", "BP"])
node.setPropertyValue("panel_field", "Sex")
# "Overlay" section
node.setPropertyValue("animation_field", "")
node.setPropertyValue("tooltip", "test")
node.setPropertyValue("normalize", True)
node.setPropertyValue("use_overlay_expr", False)
node.setPropertyValue("overlay_expression", "test")
node.setPropertyValue("records_limit", 500)
node.setPropertyValue("if_over_limit", "PlotSample")
```

表 104. multiplotnode プロパティ

multiplotnode プロパティ	データ・タイプ	プロパティの説明
x_field	フィールド	
y_fields	リスト	
panel_field	フィールド	

表 104. *multiplotnode* プロパティ (続き)

multiplotnode プロパティ	データ・タイプ	プロパティの説明
animation_field	フィールド	
normalize	フラグ	
use_overlay_expr	フラグ	
overlay_expression	<i>string</i>	
records_limit	<i>number</i>	
if_over_limit	PlotBins PlotSample PlotAll	
x_label_auto	フラグ	
x_label	<i>string</i>	
y_label_auto	フラグ	
y_label	<i>string</i>	
use_grid	フラグ	
graph_background	<i>COLOR</i>	標準のグラフ色は、このセクションの最初に説明されています。
page_background	<i>COLOR</i>	標準のグラフ色は、このセクションの最初に説明されています。

plotnode プロパティ



散布図ノードで、数値フィールド間の関係が示されます。作図は、点 (散布図) または折れ線を使用して作成できます。

例

```
node = stream.create("plot", "My node")
# "Plot" tab
node.setPropertyValue("three_D", True)
node.setPropertyValue("x_field", "BP")
node.setPropertyValue("y_field", "Cholesterol")
node.setPropertyValue("z_field", "Drug")
# "Overlay" section
node.setPropertyValue("color_field", "Drug")
node.setPropertyValue("size_field", "Age")
node.setPropertyValue("shape_field", "")
node.setPropertyValue("panel_field", "Sex")
node.setPropertyValue("animation_field", "BP")
node.setPropertyValue("transp_field", "")
node.setPropertyValue("style", "Point")
# "Output" tab
node.setPropertyValue("output_mode", "File")
node.setPropertyValue("output_format", "JPEG")
node.setPropertyValue("full_filename", "C:/temp/graph_output/
plot_output.jpeg")
```

表 105. *plotnode* プロパティ

plotnode プロパティ	データ・タイプ	プロパティの説明
x_field	フィールド	x 軸のカスタム (ユーザー設定) ラベルを指定します。ラベルでのみ使用できます。
y_field	フィールド	y 軸のカスタム (ユーザー設定) ラベルを指定します。ラベルでのみ使用できます。
three_D	フラグ	y 軸のカスタム (ユーザー設定) ラベルを指定します。3-D グラフのラベルでのみ使用できます。
z_field	フィールド	
color_field	フィールド	オーバーレイ・フィールド。
size_field	フィールド	
shape_field	フィールド	
panel_field	フィールド	各カテゴリー個別のグラフの作成に使用する名義型またはフラグ型フィールドを指定します。グラフは「パネル化」され、複数のグラフが 1 つの出力ウィンドウに表示されます。
animation_field	フィールド	アニメーションを使用して順番に表示する一連のグラフを作成してデータ値のカテゴリーを描画する、名義型またはフラグ型フィールドを指定します。
transp_field	フィールド	カテゴリーごとに異なるレベルの透過度を使用して、データ値のカテゴリーを表すフィールドを指定します。折れ線グラフでは使用できません。
overlay_type	None Smoother Function	オーバーレイ関数が表示されるか、LOESS 平滑化が表示されるかを指定します。
overlay_expression	string	overlay_type が Function に設定されているときに使用される式を指定します。
style	Point Line	

表 105. plotnode プロパティ (続き)

plotnode プロパティ	データ・タイプ	プロパティの説明
point_type	Rectangle Dot Triangle Hexagon Plus Pentagon Star BowTie HorizontalDash VerticalDash IronCross Factory House Cathedral OnionDome ConcaveTriangle OblateGlobe CatEye FourSidedPillow RoundRectangle Fan	
x_mode	Sort Overlay AsRead	
x_range_mode	Automatic UserDefined	
x_range_min	<i>number</i>	
x_range_max	<i>number</i>	
y_range_mode	Automatic UserDefined	
y_range_min	<i>number</i>	
y_range_max	<i>number</i>	
z_range_mode	Automatic UserDefined	
z_range_min	<i>number</i>	
z_range_max	<i>number</i>	
jitter	フラグ	
records_limit	<i>number</i>	
if_over_limit	PlotBins PlotSample PlotAll	

表 105. *plotnode* プロパティ (続き)

plotnode プロパティ	データ・タイプ	プロパティの説明
x_label_auto	フラグ	
x_label	<i>string</i>	
y_label_auto	フラグ	
y_label	<i>string</i>	
z_label_auto	フラグ	
z_label	<i>string</i>	
use_grid	フラグ	
graph_background	<i>COLOR</i>	標準のグラフ色は、このセクションの最初に説明されています。
page_background	<i>COLOR</i>	標準のグラフ色は、このセクションの最初に説明されています。
use_overlay_expr	フラグ	<i>overlay_type</i> の代わりに廃止される予定。

timeplotnode プロパティ



時系列グラフ・ノードで、時系列データの1つ以上のセットを表示します。通常は、最初に時間区分ノードを使用して時間ラベルフィールドを作成します。このフィールドは、X軸にラベルを付けるために使用されます。

例

```
node = stream.create("timeplot", "My node")
node.setPropertyValue("y_fields", ["sales", "men", "women"])
node.setPropertyValue("panel", True)
node.setPropertyValue("normalize", True)
node.setPropertyValue("line", True)
node.setPropertyValue("smoother", True)
node.setPropertyValue("use_records_limit", True)
node.setPropertyValue("records_limit", 2000)
# Appearance settings
node.setPropertyValue("symbol_size", 2.0)
```

表 106. *timeplotnode* プロパティ

timeplotnode プロパティ	データ・タイプ	プロパティの説明
plot_series	Series	
	Models	
use_custom_x_field	フラグ	
x_field	フィールド	
y_fields	リスト	
panel	フラグ	
normalize	フラグ	
line	フラグ	

表 106. timeplotnode プロパティ (続き)

timeplotnode プロパティ	データ・タイプ	プロパティの説明
points	フラグ	
point_type	Rectangle Dot Triangle Hexagon Plus Pentagon Star BowTie HorizontalDash VerticalDash IronCross Factory House Cathedral OnionDome ConcaveTriangle OblateGlobe CatEye FourSidedPillow RoundRectangle Fan	
smoother	フラグ	panel を True に設定した場合にのみ、平滑化を散布図に追加できます。
use_records_limit	フラグ	
records_limit	integer	
symbol_size	number	マーカー・サイズを指定します。
panel_layout	Horizontal Vertical	

eplotnode プロパティ



E-Plot (ベータ) ノードで、数値フィールド間の関係が示されます。これは散布図ノードに類似していますが、オプションは異なり、出力にはこのノードに固有の新規グラフ インターフェースを使用します。新規グラフ機能を利用するには、このベータ レベル ノードを使用します。

表 107. eplotnode プロパティ

eplotnode プロパティ	データ・タイプ	プロパティの説明
x_field	string	横の X 軸に表示するフィールドを指定します。
y_field	string	縦の Y 軸に表示するフィールドを指定します。
color_field	string	必要であれば、出力でカラー マップ オーバーレイに使用するフィールドを指定します。
size_field	string	必要であれば、出力でサイズ マップ オーバーレイに使用するフィールドを指定します。

表 107. eplotnode プロパティ (続き)

eplotnode プロパティ	データ・タイプ	プロパティの説明
shape_field	string	必要であれば、出力で形状マップオーバーレイに使用するフィールドを指定します。
interested_fields	string	出力に含めるフィールドを指定します。
records_limit	integer	出力で作図するレコードの最大数を指定します。2000 がデフォルトです。
if_over_limit	Boolean	records_limit を超えた場合に Sample オプションを使用するか、Use all data オプションを使用するかを指定します。Sample はデフォルトで、records_limit にヒットするまでデータをランダムにサンプリングします。Use all data を指定して records_limit を無視し、すべてのデータ・ポイントをプロットすると、パフォーマンスが大幅に低下する可能性があることに注意してください。

tsnnode プロパティ



t 分布 Stochastic Neighbor Embedding (t-SNE) は、高次元データの視覚化のためのツールです。t-SNE は、データ ポイントの類似性を確率に変換します。スpos モデラーの t-SNE ノードは Python で実装されており、scikit-learn® Python ライブラリーを必要とします。

表 108. tsnnode プロパティ

tsnnode プロパティ	データ・タイプ	プロパティの説明
mode_type	string	simple モードまたは expert モードを指定します。
n_components	string	埋め込み空間の次元 (2 次元 または 3 次元)。2 または 3 を指定します。デフォルトは 2 です。
method	string	barnes_hut または exact を指定します。デフォルトは barnes_hut です。
init	string	埋め込みの初期化。random または pca を指定します。デフォルトは random です。
target_field バージョン 18.2.1.1 以降は target に名前変更されています。	string	対象フィールド名。これが出力グラフのカラーマップになります。対象フィールドを指定しないと、グラフには 1 色が使用されます。

表 108. tsnenode プロパティ (続き)

tsnenode プロパティ	データ・タイプ	プロパティの説明
perplexity	Float	Perplexity は、他の多様体学習アルゴリズムで使用される最近隣の数に関連します。通常、データセットが大きいほど、必要とされる Perplexity も大きくなります。5 から 50 の間の値を選択することを考慮してください。デフォルトは 30 です。
early_exaggeration	Float	埋め込み空間における、元の空間の自然クラスターの気密度、およびクラスター間の間隔を制御します。デフォルトは 12.0 です。
learning_rate	Float	デフォルトは 200 です。
n_iter	integer	最適化の最大反復数。250 以上に設定してください。デフォルトは 1000 です。
angle	Float	あるポイントから測定した遠方ノードの角サイズ。0 から 1 の間の値を指定します。デフォルトは 0.5 です。
enable_random_seed	Boolean	random_seed パラメータを有効にするには、true に設定します。デフォルトは false です。
random_seed	integer	使用する乱数シード。デフォルトは None です。
n_iter_without_progress	integer	進捗のない最大反復数。デフォルトは 300 です。
min_grad_norm	string	勾配ノルムがこのしきい値を下回る場合、最適化は中断されます。デフォルトは 1.0E-7 です。指定できる値は次のとおりです： • 1.0E-1 • 1.0E-2 • 1.0E-3 • 1.0E-4 • 1.0E-5 • 1.0E-6 • 1.0E-7 • 1.0E-8
isGridSearch	Boolean	複数の異なる Perplexity で t-SNE を実行するには、true に設定します。デフォルトは false です。
output_Rename	Boolean	カスタム名を設定する場合は true を指定し、出力に名前を自動的に付ける場合は false を指定します。デフォルトは false です。

表 108. tsnode プロパティ (続き)

tsnode プロパティ	データ・タイプ	プロパティの説明
output_to	string	Screen または Output を指定します。デフォルトは Screen です。
full_filename	string	出力ファイル名を指定します。
output_file_type	string	出力ファイル形式。HTML または Output object を指定します。デフォルトは HTML です。

webnode プロパティ



Web グラフ・ノードで、複数のシンボル値(カテゴリー) フィールドの値の関係の強さが示されます。このグラフでは、接続の強さを示すためにさまざまな幅の線が使用されます。Web グラフ・ノードを使用すると、例えば、E コマース・サイトでの一連の商品の購入間の関係を調査できます。

例

```
node = stream.create("web", "My node")
# "Plot" tab
node.setPropertyValue("use_directed_web", True)
node.setPropertyValue("to_field", "Drug")
node.setPropertyValue("fields", ["BP", "Cholesterol", "Sex", "Drug"])
node.setPropertyValue("from_fields", ["BP", "Cholesterol", "Sex"])
node.setPropertyValue("true_flags_only", False)
node.setPropertyValue("line_values", "Absolute")
node.setPropertyValue("strong_links_heavier", True)
# "Options" tab
node.setPropertyValue("max_num_links", 300)
node.setPropertyValue("links_above", 10)
node.setPropertyValue("num_links", "ShowAll")
node.setPropertyValue("discard_links_min", True)
node.setPropertyValue("links_min_records", 5)
node.setPropertyValue("discard_links_max", True)
node.setPropertyValue("weak_below", 10)
node.setPropertyValue("strong_above", 19)
node.setPropertyValue("link_size_continuous", True)
node.setPropertyValue("web_display", "Circular")
```

表 109. webnode プロパティ

webnode プロパティ	データ・タイプ	プロパティの説明
use_directed_web	フラグ	
fields	リスト	
to_field	フィールド	
from_fields	リスト	
true_flags_only	フラグ	

表 109. webnode プロパティ (続き)

webnode プロパティ	データ・タイプ	プロパティの説明
line_values	Absolute OverallPct PctLarger PctSmaller	
strong_links_heavier	フラグ	
num_links	ShowMaximum ShowLinksAbove ShowAll	
max_num_links	number	
links_above	number	
discard_links_min	フラグ	
links_min_records	number	
discard_links_max	フラグ	
links_max_records	number	
weak_below	number	
strong_above	number	
link_size_continuous	フラグ	
web_display	Circular Network Directed Grid	
graph_background	COLOR	標準のグラフ色は、このセクションの最初に説明されています。
symbol_size	number	マーカー・サイズを指定します。

第 13 章 モデル作成ノードのプロパティ

一般的なモデル作成ノードのプロパティ

次のプロパティは、複数またはすべてのモデル作成ノードに共通です。個別のモデル作成ノードに関しては、必要に応じてドキュメント内に例外を記載しています。

表 110. 一般的なモデル作成ノードのプロパティ		
プロパティ	値	プロパティの説明
custom_fields	フラグ	真 (true) の場合は、現在のノードのターゲット、入力、その他フィールドなどを指定することができます。偽 (false) の場合は、上流のデータ型ノードから現在の設定が使用されます。
target または targets	フィールド または [field1 ... fieldN]	モデルのタイプによって、単一の対象フィールドまたは複数の対象フィールドを指定します。
inputs	[field1 ... fieldN]	モデルで使用される入力または予測変数フィールド。
partition	フィールド	
use_partitioned_data	フラグ	区分フィールドが定義される場合、このオプションは学習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。
use_split_data	フラグ	
splits	[field1 ... fieldN]	分割モデル作成に使用する、フィールドを選択します。use_split_data が True に設定されている場合にのみ有効です。
use_frequency	フラグ	各モデル・タイプで言及するとおり、重みフィールドおよび度数フィールドが特定のモデルで使用されます。
frequency_field	フィールド	
use_weight	フラグ	
weight_field	フィールド	
use_model_name	フラグ	
model_name	string	ユーザーが指定する新規モデル名。
mode	Simple Expert	

anomalydetectionnode プロパティ



異常値検査ノードで、「正常な」データのパターンに合致しない異常ケースや外れ値を識別します。このノードを使用すると、既知のパターンに適合しない場合や、探している内容が正確にわからない場合でも、外れ値を識別することができます。

例

```
node = stream.create("anomalydetection", "My node")
node.setPropertyValue("anomaly_method", "PerRecords")
node.setPropertyValue("percent_records", 95)
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("peer_group_num_auto", True)
node.setPropertyValue("min_num_peer_groups", 3)
node.setPropertyValue("max_num_peer_groups", 10)
```

表 111. anomalydetectionnode プロパティ

anomalydetectionnode プロパティ	値	プロパティの説明
inputs	[field1 ... fieldN]	異常値検査モデルは、指定の入力フィールドに基づいてレコードをスクリーニングします。ターゲット・フィールドは使用しません。重みフィールドおよび度数フィールドも使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
mode	Expert Simple	
anomaly_method	IndexLevel PerRecords NumRecords	レコードに異常としてフラグを設定するための、分割値を決めるのに使用される方法を指定します。
index_level	number	異常としてフラグを設定するための最小分割値を指定します。
percent_records	number	学習データ内のレコードの割合 (%) に基づいてレコードにフラグを設定するための、閾値を設定します。
num_records	number	学習データ内のレコードの数に基づいてレコードにフラグを設定するための、閾値を設定します。
num_fields	integer	各異常レコードに報告するフィールド数。
impute_missing_values	フラグ	
adjustment_coeff	number	距離の計算時に連続型とカテゴリ・フィールド間に指定された関連の重みのバランスをとるために使用される値。

表 111. *anomalydetectionnode* プロパティ (続き)

anomalydetectionnode プロパティ	値	プロパティの説明
peer_group_num_auto	フラグ	ピア・グループ数を自動的に計算します。
min_num_peer_groups	<i>integer</i>	peer_group_num_auto が True に設定されている場合に使用されるピア・グループの最小数を指定します。
max_num_per_groups	<i>integer</i>	ピア・グループの最大数を指定します。
num_peer_groups	<i>integer</i>	peer_group_num_auto が False に設定されている場合に使用されるピア・グループの数を指定します。
noise_level	<i>number</i>	クラスタリング中の外れ値の処理方法を決定します。0 から 0.5 までの値を指定してください。
noise_ratio	<i>number</i>	ノイズのバッファリングに使用されるコンポーネントに割り当てられる、メモリーの量を指定します。0 から 0.5 までの値を指定してください。

apriorinode プロパティ



Apriori ノードで、データからルール・セットを抽出し、情報内容が最も充実したルールを引き出します。Apriori には、5 種類のルール選択方法があり、高度なインデックス作成方法を使用して、大きなデータ・セットが効率的に処理されます。大きな問題の場合は、一般に、Apriori の方が高速に学習できます。保持できるルール数に特に制限はありません。また、最大 32 の前提条件を持つルールを処理できます。Apriori は、入力フィールドと出力フィールドをすべてカテゴリー化する必要がありますが、このタイプのデータ用に最適化されているため、パフォーマンスが向上します。

例

```
node = stream.create("apriori", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("partition", "Test")
# For non-transactional
node.setPropertyValue("use_transactional_data", False)
node.setPropertyValue("consequents", ["Age"])
node.setPropertyValue("antecedents", ["BP", "Cholesterol", "Drug"])
# For transactional
node.setPropertyValue("use_transactional_data", True)
node.setPropertyValue("id_field", "Age")
node.setPropertyValue("contiguous", True)
node.setPropertyValue("content_field", "Drug")
# "Model" tab
node.setPropertyValue("use_model_name", False)
node.setPropertyValue("model_name", "Apriori_bp_choles_drug")
node.setPropertyValue("min_supp", 7.0)
node.setPropertyValue("min_conf", 30.0)
node.setPropertyValue("max_antecedents", 7)
node.setPropertyValue("true_flags", False)
node.setPropertyValue("optimize", "Memory")
```

```
# "Expert" tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("evaluation", "ConfidenceRatio")
node.setPropertyValue("lower_bound", 7)
```

表 112. *apriorinode* プロパティ

apriorinode プロパティ	値	プロパティの説明
consequents	フィールド	Apriori モデルは標準的な対象フィールドおよび入力フィールドの結果と条件を使用します。重みフィールドおよび度数フィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
antecedents	[field1 ... fieldN]	
min_supp	number	
min_conf	number	
max_antecedents	number	
true_flags	フラグ	
optimize	Speed Memory	
use_transactional_data	フラグ	値が true である場合、各トランザクション ID のスコアは他のトランザクション ID に依存しません。スコアリング対象のデータが大量過ぎて、許容し得るパフォーマンスが得られない場合は、データを分離することをお勧めします。
contiguous	フラグ	
id_field	string	
content_field	string	
mode	Simple Expert	
evaluation	RuleConfidence DifferenceToPrior ConfidenceRatio InformationDifference NormalizedChiSquare	
lower_bound	number	

表 112. *apriorinode* プロパティ (続き)

apriorinode プロパティ	値	プロパティの説明
optimize	Speed Memory	モデル作成が速度とメモリーのどちらにより最適化されるかを指定します。

associationrulesnode プロパティ



アソシエーション・ルール・ノードは Apriori ノードに似ていますが、Apriori とは異なり、アソシエーション・ルール・ノードはリスト・データを処理できます。さらに、アソシエーション・ルール・ノードを IBMSPSS Analytic Server と共に使用すると、ビッグデータの処理や高速な並列処理の利用が可能になります。

表 113. *associationrulesnode* プロパティ

associationrulesnode プロパティ	データ・タイプ	プロパティの説明
predictions	フィールド	このリスト内の各フィールドは、ルールの予測フィールドとしてのみ表示することができます。
conditions	[field1...fieldN]	このリスト内の各フィールドは、ルールの条件としてのみ表示することができます。
max_rule_conditions	integer	1 つのルールに含めることができる条件の最大数。最小値は 1、最大値は 9 です。
max_rule_predictions	integer	1 つのルールに含めることができる予測の最大数。最小値は 1、最大値は 5 です。
max_num_rules	integer	ルール構築の一部としてみなすことができるルールの最大数。最小値は 1、最大値は 10,000 です。
rule_criterion_top_n	Confidence Rulesupport Lift Conditionsupport Deployability	値を判断するルール基準。この基準により、モデル内の上位 N 件のルールが選択されます。
true_flags	Boolean	これを Y に設定すると、ルールの構築時に、true の値を持つフラグ フィールドだけが処理対象になります。
rule_criterion	Boolean	これを Y に設定すると、モデルの構築時に、ルール基準の値を使用してルールが除外されます。

表 113. associationrulesnode プロパティ (続き)

associationrulesnode プロパティ	データ・タイプ	プロパティの説明
min_confidence	number	0.1 から 100: モデルによって生成されたルールについて最低限必要な確信度レベルのパーセント値。ここで指定された値よりも低い確信度レベルを持つルールがモデルによって生成された場合、そのルールは破棄されます。
min_rule_support	number	0.1 から 100: モデルによって生成されたルールについて最低限必要なルール サポートのパーセント値。ここで指定された値よりも低いルール サポート レベルを持つルールがモデルによって生成された場合、そのルールは破棄されます。
min_condition_support	number	0.1 から 100: モデルによって生成されたルールについて最低限必要な条件サポートのパーセント値。ここで指定された値よりも低い条件サポート レベルを持つルールがモデルによって生成された場合、そのルールは破棄されます。
min_lift	integer	1 から 10: モデルによって生成されたルールについて最低限必要なリフト レベルを表します。ここで指定された値よりも低いリフト レベルを持つルールがモデルによって生成された場合、そのルールは破棄されます。
exclude_rules	Boolean	このプロパティを使用して、モデルによるルールの作成元として使用しない関連フィールドのリストを選択します。 例: <code>set :gsarsnode.exclude_rules = [[field1,field2, field3]], [[field4, field5]]-[]</code> で区切られたフィールドの各リストは、テーブル内の行です。
num_bins	integer	連続型フィールドのビン分割先となる自動ビンの数を設定します。最小値は 2、最大値は 10 です。
max_list_length	integer	最大長が不明なすべてのリスト フィールドに適用されます。ここで指定された数を上限として、リスト内の要素がモデルの構築で使用されます。ここで指定された数を超える要素については、すべて破棄されます。最小値は 1、最大値は 100 です。
output_confidence	Boolean	
output_rule_support	Boolean	
output_lift	Boolean	
output_condition_support	Boolean	
output_deployability	Boolean	

表 113. associationrulesnode プロパティ (続き)

associationrulesnode プロパティ	データ・タイプ	プロパティの説明
rules_to_display	upto all	出力テーブルに表示されるルール of 最大数。
display_upto	integer	rules_to_display で upto を設定した場合は、出力テーブルに表示されるルール of 数を指定します。最小値は 1 です。
field_transformations	Boolean	
records_summary	Boolean	
rule_statistics	Boolean	
most_frequent_values	Boolean	
most_frequent_fields	Boolean	
word_cloud	Boolean	
word_cloud_sort	Confidence Rulesupport Lift Conditionsupport Deployability	
word_cloud_display	integer	最小値は 1、最大値は 20 です。
max_predictions	integer	スコアに対する各入力に適用できるルール of 最大数。
criterion	Confidence Rulesupport Lift Conditionsupport Deployability	ルール of 強度を判断するための尺度を選択します。
allow_repeats	Boolean	同じ予測を持つルール of スコア内に含めるかどうかを決定します。
check_input	NoPredictions Predictions NoCheck	

autoclassifiernode プロパティ



自動分類ノードは、2 種類の結果 (yes/no、 churn/don't churn など) を生じる多くの異なるモデルを作成および比較し、与えられた分析への最善のアプローチを選ぶことができるようになります。多くのモデル作成アルゴリズムに対応し、希望する方法、各特定のオプション、そして結果を比較するための基準を選択することができます。ノードは、指定されたオプションに基づいてモデルのセットを生成し、指定された基準に従って最適な候補をランク付けします。

例

```
node = stream.create("autoclassifier", "My node")
node.setPropertyValue("ranking_measure", "Accuracy")
node.setPropertyValue("ranking_dataset", "Training")
node.setPropertyValue("enable_accuracy_limit", True)
node.setPropertyValue("accuracy_limit", 0.9)
node.setPropertyValue("calculate_variable_importance", True)
node.setPropertyValue("use_costs", True)
node.setPropertyValue("svm", False)
```

表 114. autoclassifiernode プロパティ

autoclassifiernode プロパティ	値	プロパティの説明
target	フィールド	フラグ型対照の場合、自動分類ノードは 1 つの対象フィールドおよび 1 つ以上の入力フィールドを使用します。重みフィールドおよび度数フィールドも指定することができます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
ranking_measure	Accuracy Area_under_curve Profit Lift Num_variables	
ranking_dataset	Training Test	
number_of_models	integer	モデル・ナゲットに含まれるモデルの数。1 と 100 の間の整数を指定します。
calculate_variable_importance	フラグ	
enable_accuracy_limit	フラグ	
accuracy_limit	integer	0 と 100 の間の整数です。

表 114. autoclassifiernode プロパティ (続き)

autoclassifiernode プロパティ	値	プロパティの説明
enable_area_under_curve_limit	フラグ	
area_under_curve_limit	number	0.0 と 1.0 の間の実数。
enable_profit_limit	フラグ	
profit_limit	number	1 以上の整数。
enable_lift_limit	フラグ	
lift_limit	number	1.0 を超える実数。
enable_number_of_variables_limit	フラグ	
number_of_variables_limit	number	1 以上の整数。
use_fixed_cost	フラグ	
fixed_cost	number	0.0 を超える実数。
variable_cost	フィールド	
use_fixed_revenue	フラグ	
fixed_revenue	number	0.0 を超える実数。
variable_revenue	フィールド	
use_fixed_weight	フラグ	
fixed_weight	number	0.0 を超える実数。
variable_weight	フィールド	
lift_percentile	number	0 と 100 の間の整数です。
enable_model_build_time_limit	フラグ	
model_build_time_limit	number	個々のモデルのそれぞれを構築するためにかかる時間を制限するために分数を設定する整数。
enable_stop_after_time_limit	フラグ	
stop_after_time_limit	number	自動分類の実行のための全体経過時間を制限するために時間数を設定する実数。
enable_stop_after_valid_model_produced	フラグ	
use_costs	フラグ	
<algorithm>	フラグ	特定のアルゴリズムの使用の有効、無効を切り替えます。
<algorithm>.<property>	string	特定のアルゴリズムのプロパティ値を設定します。詳しくは、トピック 226 ページの『 アルゴリズム・プロパティの設定 』を参照してください。

アルゴリズム・プロパティの設定

自動分類ノード、自動数値ノード、自動クラスタリング・ノードについては、ノードが使用する特定のアルゴリズムのプロパティは、次の一般形式を使用して設定できます。

```
autonode.setKeyedPropertyValue(<algorithm>, <property>, <value>)
```

以下に例を示します。

```
node.setKeyedPropertyValue("neuralnetwork", "method", "MultilayerPerceptron")
```

自動分類ノードのアルゴリズム名は、cart、chaid、quest、c50、logreg、decisionlist、bayesnet、discriminant、svm および knn です。

自動数値ノードのアルゴリズム名は、cart、chaid、neuralnetwork、genlin、svm、regression、linear および knn です。

自動クラスタリング・ノードのアルゴリズム名は、twostep、k-means、および kohonen です。

プロパティ名は、各アルゴリズムノードのために文書化されている標準です。

ピリオドなどの句読点を含むアルゴリズム・プロパティは、次のように一重引用符で囲む必要があります。

```
node.setKeyedPropertyValue("logreg", "tolerance", "1.0E-5")
```

次のように、複数の値をプロパティに割り当てることもできます。

```
node.setKeyedPropertyValue("decisionlist", "search_direction", ["Up",  
"Down"])
```

特定のアルゴリズムの使用の有効、無効を切り替えるには、次のようにします。

```
node.setPropertyValue("chaid", True)
```

注: 自動分類ノードで特定のアルゴリズム・オプションが使用可能でない場合、または値の範囲ではなく、1つの値だけを指定できるときは、標準の方法でノードにアクセスするときと同じ制限が、スクリプトにも適用されます。

autoclusternode プロパティ



自動クラスタリング・ノードは、同様の特性を持つレコードのグループを識別するクラスタリング・モデルを推定し、比較します。ノードは他の自動化モデル作成ノードと同じように動作し、複数の組み合わせのオプションを単一のモデル作成の実行で検証できます。モデルは、クラスター・モデルの有用性をフィルタリングおよびランク付けする基本的な指標を使用して比較し、特定のフィールドの重要度に基づいて指標を提供します。

例

```
node = stream.create("autocluster", "My node")  
node.setPropertyValue("ranking_measure", "Silhouette")  
node.setPropertyValue("ranking_dataset", "Training")  
node.setPropertyValue("enable_silhouette_limit", True)  
node.setPropertyValue("silhouette_limit", 5)
```

表 115. autoclusternode プロパティ

autoclusternode プロパティ	値	プロパティの説明
evaluation	フィールド	注: のみ。重要度の値を計算するフィールドを識別します。また、どれだけクラスターがフィールドの値を区別するか、どれだけ正確にモデルがこのフィールドを予測するかを識別するために使用することができます。
ranking_measure	Silhouette Num_clusters Size_smallest_cluster Size_largest_cluster Smallest_to_largest Importance	
ranking_dataset	Training Test	
summary_limit	integer	レポートに一覧するモデルの数。1 と 100 の間の整数を指定します。
enable_silhouette_limit	フラグ	
silhouette_limit	integer	0 と 100 の間の整数です。
enable_number_less_limit	フラグ	
number_less_limit	number	0.0 と 1.0 の間の実数。
enable_number_greater_limit	フラグ	
number_greater_limit	number	1 以上の整数。
enable_smallest_cluster_limit	フラグ	
smallest_cluster_units	Percentage Counts	
smallest_cluster_limit_percentage	number	
smallest_cluster_limit_count	integer	1 以上の整数。
enable_largest_cluster_limit	フラグ	

表 115. autoclusternode プロパティ (続き)

autoclusternode プロパティ	値	プロパティの説明
largest_cluster_units	Percentage Counts	
largest_cluster_limit_percentage	number	
largest_cluster_limit_count	integer	
enable_smallest_largest_limit	フラグ	
smallest_largest_limit	number	
enable_importance_limit	フラグ	
importance_limit_condition	Greater_than Less_than	
importance_limit_greater_than	number	0 と 100 の間の整数です。
importance_limit_less_than	number	0 と 100 の間の整数です。
<algorithm>	フラグ	特定のアルゴリズムの使用の有効、無効を切り替えます。
<algorithm>.<property>	string	特定のアルゴリズムのプロパティ値を設定します。詳しくは、トピック 226 ページの『アルゴリズム・プロパティの設定』を参照してください。

autonumericnode プロパティ



自動数値ノードでは、多くのさまざまな方法を使用し、連続する数値範囲の結果を求めてモデルを推定し比較します。このノードは、自動分類ノードと同じ方法で動作し、1回のモデル作成のパスで、複数の組み合わせのオプションを使用し試すアルゴリズムを選択することができます。使用できるアルゴリズムには、ニューラル・ネットワーク、C&R Tree、CHAID、線型、一般化線型、サポート・ベクトル・マシン (SVM) が含まれています。モデルは、相関、相対誤差、または使用される変数の数に基づいて比較できます。

例

```
node = stream.create("autonumeric", "My node")
node.setPropertyValue("ranking_measure", "Correlation")
node.setPropertyValue("ranking_dataset", "Training")
node.setPropertyValue("enable_correlation_limit", True)
node.setPropertyValue("correlation_limit", 0.8)
node.setPropertyValue("calculate_variable_importance", True)
node.setPropertyValue("neuralnetwork", True)
node.setPropertyValue("chaid", False)
```

表 116. autonumericnode プロパティ

autonumericnode プロパティ	値	プロパティの説明
custom_fields	フラグ	真 (True) の場合、データ型ノード設定の代わりにカスタム・フィールド設定が使用されます。
target	フィールド	自動数値ノードは 1 つの対象フィールドおよび 1 つ以上の入力フィールドを使用します。重みフィールドおよび度数フィールドも指定することができます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
inputs	[field1 ... field2]	
partition	フィールド	
use_frequency	フラグ	
frequency_field	フィールド	
use_weight	フラグ	
weight_field	フィールド	
use_partitioned_data	フラグ	データ区分フィールドが定義されている場合、学習データだけがモデルの構築に使用されます。
ranking_measure	Correlation NumberOfFields	
ranking_dataset	Test Training	
number_of_models	integer	モデル・ナゲットに含まれるモデルの数。1 と 100 の間の整数を指定します。
calculate_variable_importance	フラグ	
enable_correlation_limit	フラグ	
correlation_limit	integer	
enable_number_of_fields_limit	フラグ	
number_of_fields_limit	integer	
enable_relative_error_limit	フラグ	
relative_error_limit	integer	
enable_model_build_time_limit	フラグ	

表 116. *autonumericnode* プロパティ (続き)

autonumericnode プロパティ	値	プロパティの説明
model_build_time_limit	<i>integer</i>	
enable_stop_after_time_limit	フラグ	
stop_after_time_limit	<i>integer</i>	
stop_if_valid_model	フラグ	
<algorithm>	フラグ	特定のアルゴリズムの使用の有効、無効を切り替えます。
<algorithm>.<property>	<i>string</i>	特定のアルゴリズムのプロパティ値を設定します。詳しくは、トピック 226 ページの『 アルゴリズム・プロパティの設定 』を参照してください。

bayesnetnode プロパティ



Bayesian network (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、発生の尤度を確立できます。このノードは、主に分類に使用される Tree Augmented Naive Bayes (TAN) および Markov Blanket ネットワークに焦点を当てています。

例

```
node = stream.create("bayesnet", "My node")
node.setPropertyValue("continue_training_existing_model", True)
node.setPropertyValue("structure_type", "MarkovBlanket")
node.setPropertyValue("use_feature_selection", True)
# Expert tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("all_probabilities", True)
node.setPropertyValue("independence", "Pearson")
```

表 117. *bayesnetnode* プロパティ

bayesnetnode プロパティ	値	プロパティの説明
inputs	[<i>field1 ... fieldN</i>]	Bayesian network (ベイズ) モデルは単一の対象フィールドおよび 1 つ以上の入力フィールドを使用します。連続フィールドは自動的に分割されます。詳しくは、トピック 217 ページの『 一般的なモデル作成ノードのプロパティ 』を参照してください。
continue_training_existing_model	フラグ	
structure_type	TAN MarkovBlanket	Bayesian Network (ベイズ) を構築時に使用する構造を選択します。
use_feature_selection	フラグ	

表 117. bayesnetnode プロパティ (続き)

bayesnetnode プロパティ	値	プロパティの説明
parameter_learning_method	Likelihood Bayes	親の値が認識されるノード間の条件付き確率テーブルを推定するために用いる方法を指定します。
mode	Expert Simple	
missing_values	フラグ	
all_probabilities	フラグ	
independence	Likelihood Pearson	2 つの変数のペアの観測が互いに独立しているかどうかを評価するために用いる方法を指定します。
significance_level	number	独立性を判断するための分割値を指定します。
maximal_conditioning_set	number	独立性検定に使用する条件変数の最大数を指定します。
inputs_always_selected	[field1 ... fieldN]	Bayesian network (ベイズ) 構築時にデータセットのどのフィールドを常に使用するかを指定します。 注: 対象フィールドは必ず選択されます。
maximum_number_inputs	number	Bayesian network (ベイズ) 構築で使用する入力フィールドの最大数を指定します。
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

buildr プロパティ



R 構築ノードを使用すると、IBM スポス モデラー に展開されているモデル作成およびモデル・スコアリングを実行するためのカスタムの R スクリプトを入力できます。

例

```
node = stream.create("buildr", "My node")
node.setPropertyValue("score_syntax", "")
result<-predict(modelerModel,newdata=modelerData)
```

```

modelerData<-cbind(modelerData,result)
var1<-
c(fieldName="NaPrediction",fieldLabel="",fieldStorage="real",fieldMeasure="",
fieldFormat="",fieldRole="")
modelerDataModel<-data.frame(modelerDataModel,var1)""")

```

表 118. buildr プロパティ

buildr プロパティ	値	プロパティの説明
build_syntax	string	モデル作成用の R スクリプト・シンタックス。
score_syntax	string	モデル・スコアリング用の R スクリプト・シンタックス。
convert_flags	StringsAndDoubles LogicalValues	フラグ型フィールドを変換するためのオプション。
convert_datetime	フラグ	日付形式または日付/時刻形式の変数を R の日付/時刻形式に変換するためのオプション。
convert_datetime_class	POSIXct POSIXlt	日付形式または日付/時刻形式の変数のうち、どの形式の変数を変換するかを指定するためのオプション。
convert_missing	フラグ	欠損値を R NA 値に変換するオプションです。
output_html	フラグ	R モデル・ナゲットのタブにグラフを表示するためのオプション。
output_text	フラグ	R モデル・ナゲットのタブに R コンソールのテキスト出力を書き込むためのオプション。

c50node プロパティ



C5.0 ノードは、ディシジョン・ツリーとルール・セットのどちらかを構築します。このモデルは、各レベルで最大の情報の対応をもたらすフィールドに基づいてサンプルを分割します。対象フィールドは、カテゴリーでなければなりません。複数の分割を 2 つ以上のサブグループに分割できます。

例

```

node = stream.create("c50", "My node")
# "Model" tab
node.setPropertyValue("use_model_name", False)
node.setPropertyValue("model_name", "C5_Drug")
node.setPropertyValue("use_partitioned_data", True)
node.setPropertyValue("output_type", "DecisionTree")
node.setPropertyValue("use_xval", True)
node.setPropertyValue("xval_num_folds", 3)
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("favor", "Generality")
node.setPropertyValue("min_child_records", 3)
# "Costs" tab
node.setPropertyValue("use_costs", True)
node.setPropertyValue("costs", [["drugA", "drugX", 2]])

```

表 119. c50node プロパティ

c50node プロパティ	値	プロパティの説明
target	フィールド	C50 モデルは単一の対象フィールドおよび 1 つ以上の入力フィールドを使用します。重みフィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
output_type	DecisionTree RuleSet	
group_symbolics	フラグ	
use_boost	フラグ	
boost_num_trials	number	
use_xval	フラグ	
xval_num_folds	number	
mode	Simple Expert	
favor	Accuracy Generality	精度 (Accuracy) または一般化 (Generality) を選択。
expected_noise	number	
min_child_records	number	
pruning_severity	number	
use_costs	フラグ	
costs	構造化	これは構造化されたプロパティです。
use_winnowing	フラグ	
use_global_pruning	フラグ	デフォルトではオン (True)。
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

carmanode プロパティ



CARMA モデルは、入力または対象フィールドを指定しなくても、データからルールのセットを抽出します。</all>and GRI, END FILTER ALL -->Apriori とは対照的に、CARMA ノードでは、前提条件サポートだけではなく、ルール・サポート (前提条件と結果の両方のサポート) を対象とした構築の設定が可能です。これは、生成されたルールをさまざまなアプリケーションで活用できることを意味します。例えば、この休暇シーズンに販売促進する項目を結果とする、商品またはサービス (前提条件) のリストを調べることができます。

例

```
node = stream.create("carma", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("use_transactional_data", True)
node.setPropertyValue("inputs", ["BP", "Cholesterol", "Drug"])
node.setPropertyValue("partition", "Test")
# "Model" tab
node.setPropertyValue("use_model_name", False)
node.setPropertyValue("model_name", "age_bp_drug")
node.setPropertyValue("use_partitioned_data", False)
node.setPropertyValue("min_supp", 10.0)
node.setPropertyValue("min_conf", 30.0)
node.setPropertyValue("max_size", 5)
# Expert Options
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("use_pruning", True)
node.setPropertyValue("pruning_value", 300)
node.setPropertyValue("vary_support", True)
node.setPropertyValue("estimated_transactions", 30)
node.setPropertyValue("rules_without_antecedents", True)
```

表 120. carmanode プロパティ		
carmanode プロパティ	値	プロパティの説明
inputs	[field1 ... fieldn]	CARMA モデルは対象フィールドでなく、入力フィールドのリストを使用します。 重みフィールドおよび度数フィールドは使用しません。 詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
id_field	フィールド	モデル作成の ID フィールドとして使用するフィールド。
contiguous	フラグ	ID フィールドの ID が連続するかどうかを指定します。
use_transactional_data	フラグ	
content_field	フィールド	
min_supp	数値 (パーセント)	前提条件範囲(サポート)ではなく、ルール範囲に関連します。 デフォルト値は 20% です。
min_conf	数値 (パーセント)	デフォルト値は 20% です。
max_size	number	デフォルト値は 10 です。

表 120. carmanode プロパティ (続き)

carmanode プロパティ	値	プロパティの説明
mode	Simple Expert	デフォルトは Simple です。
exclude_multiple	フラグ	複数の結果を持つルールを除外します。デフォルトは False です。
use_pruning	フラグ	デフォルトは False です。
pruning_value	number	デフォルトは 500 です。
vary_support	フラグ	
estimated_transactions	integer	
rules_without_antecedents	フラグ	

cartnode プロパティ



C&R Tree (分類と回帰ツリー) ノードは、ディシジョン・ツリーを生成し、将来の観測値を予測または分類できるようにします。この方法は再帰的なデータ区分を使用して学習レコードを複数のセグメントに分割し、各ステップで不純性を最小限に抑えます。ツリーのノードが「純粹」であると考えられるのは、ノード中にあるケースの 100% が、対象フィールドのある特定のカテゴリに分類される場合です。対象フィールドと入力フィールドには、数値範囲またはカテゴリ (名義型、順序型、またはフラグ型) を指定できます。すべての分割は 2 値 (2 つのサブグループのみ) です。

例

```
node = stream.createAt("cart", "My node", 200, 100)
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("target", "Drug")
node.setPropertyValue("inputs", ["Age", "BP", "Cholesterol"])
# "Build Options" tab, "Objective" panel
node.setPropertyValue("model_output_type", "InteractiveBuilder")
node.setPropertyValue("use_tree_directives", True)
node.setPropertyValue("tree_directives", "" "Grow Node Index 0 Children 1 2
Grow Node Index 2 Children 3 4""")
# "Build Options" tab, "Basics" panel
node.setPropertyValue("prune_tree", False)
node.setPropertyValue("use_std_err_rule", True)
node.setPropertyValue("std_err_multiplier", 3.0)
node.setPropertyValue("max_surrogates", 7)
# "Build Options" tab, "Stopping Rules" panel
node.setPropertyValue("use_percentage", True)
node.setPropertyValue("min_parent_records_pc", 5)
node.setPropertyValue("min_child_records_pc", 3)
# "Build Options" tab, "Advanced" panel
node.setPropertyValue("min_impurity", 0.0003)
node.setPropertyValue("impurity_measure", "Twoing")
# "Model Options" tab
node.setPropertyValue("use_model_name", True)
node.setPropertyValue("model_name", "Cart_Drug")
```

表 121. cartnode プロパティ

cartnode プロパティ	値	プロパティの説明
target	フィールド	C&R Tree モデルは 1 つの対象フィールドおよび 1 つ以上の入力フィールドを使用します。度数フィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
continue_training_existing_model	フラグ	
objective	Standard Boosting Bagging psm	psm は非常に大きいデータセットに使用され、Server の接続が必要です。
model_output_type	Single InteractiveBuilder	
use_tree_directives	フラグ	
tree_directives	string	ツリーの成長のためのディレクティブ (式) を指定します。ディレクティブ (式) は、改行や引用符のエスケープ処理を回避するために、三重の引用符で囲むことができます。ディレクティブは、データやモデルリング・オプションの些細な変更依存するため、他のデータセットに対しては一般化できません。
use_max_depth	Default Custom	
max_depth	integer	最大ツリー深さ (0 から 1000)。use_max_depth = Custom の場合にのみ使用されます。
prune_tree	フラグ	オーバーフィットしないようにツリーを剪定します。
use_std_err	フラグ	リスクにおける最大差 (標準誤差) を使用します。
std_err_multiplier	number	最大差。
max_surrogates	number	最大代理変数。
use_percentage	フラグ	
min_parent_records_pc	number	
min_child_records_pc	number	

表 121. cartnode プロパティ (続き)

cartnode プロパティ	値	プロパティの説明
min_parent_records_abs	<i>number</i>	
min_child_records_abs	<i>number</i>	
use_costs	フラグ	
costs	構造化	構造化プロパティ。
priors	Data Equal Custom	
custom_priors	構造化	構造化プロパティ。
adjust_priors	フラグ	
trails	<i>number</i>	ブーストまたはバグのコンポーネント・モデル数。
set_ensemble_method	Voting HighestProbability HighestMeanProbability	カテゴリー型対象のデフォルト結合ルール。
range_ensemble_method	Mean Median	連続型対象のデフォルト結合ルール。
large_boost	フラグ	特に大きなデータセットのブースティングを適用します。
min_impurity	<i>number</i>	
impurity_measure	Gini Twoing Ordered	
train_pct	<i>number</i>	オーバーフィット防止セット。
set_random_seed	フラグ	結果を再現オプション。
seed	<i>number</i>	
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

表 121. cartnode プロパティ (続き)

cartnode プロパティ	値	プロパティの説明
adjusted_propensity_partition	Test Validation	

chaidnode プロパティ



CHAID ノードはディシジョン・ツリーを生成し、カイ二乗統計値を使用して最適な分割を識別します。C&R Tree および QUEST ノードと異なり、CHAID は、非 2 分岐ツリーを生成できます。これは、ある分岐が 3 個以上のブランチを持てることを意味します。対象フィールドおよび入力フィールドは、数値範囲 (連続型) またはカテゴリとなります。拡張 CHAID は、CHAID を変更したものです。これにより、考えられるすべての分割をより徹底的に調べることができますが、計算に時間がかかります。

例

```

filenode = stream.createAt("variablefile", "My node", 100, 100)
filenode.setPropertyValue("full_filename", "$CLEO_DEMOS/DRUG1n")
node = stream.createAt("chaid", "My node", 200, 100)
stream.link(filenode, node)

node.setPropertyValue("custom_fields", True)
node.setPropertyValue("target", "Drug")
node.setPropertyValue("inputs", ["Age", "Na", "K", "Cholesterol", "BP"])
node.setPropertyValue("use_model_name", True)
node.setPropertyValue("model_name", "CHAID")
node.setPropertyValue("method", "Chaid")
node.setPropertyValue("model_output_type", "InteractiveBuilder")
node.setPropertyValue("use_tree_directives", True)
node.setPropertyValue("tree_directives", "Test")
node.setPropertyValue("split_alpha", 0.03)
node.setPropertyValue("merge_alpha", 0.04)
node.setPropertyValue("chi_square", "Pearson")
node.setPropertyValue("use_percentage", False)
node.setPropertyValue("min_parent_records_abs", 40)
node.setPropertyValue("min_child_records_abs", 30)
node.setPropertyValue("epsilon", 0.003)
node.setPropertyValue("max_iterations", 75)
node.setPropertyValue("split_merged_categories", True)
node.setPropertyValue("bonferroni_adjustment", True)

```

表 122. chaidnode プロパティ

chaidnode プロパティ	値	プロパティの説明
target	フィールド	CHAID モデルは単一の対象フィールドおよび 1 つ以上の入力フィールドを使用します。度数フィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
continue_training_existing_model	フラグ	

表 122. chaidnode プロパティ (続き)

chaidnode プロパティ	値	プロパティの説明
objective	Standard Boosting Bagging psm	psm は非常に大きいデータセットに使用され、Server の接続が必要です。
model_output_type	Single InteractiveBuilder	
use_tree_directives	フラグ	
tree_directives	string	
method	Chaid ExhaustiveChaid	
use_max_depth	Default Custom	
max_depth	integer	最大ツリー深さ (0 から 1000)。 use_max_depth = Custom の場合にのみ使用されます。
use_percentage	フラグ	
min_parent_records_pc	number	
min_child_records_pc	number	
min_parent_records_abs	number	
min_child_records_abs	number	
use_costs	フラグ	
costs	構造化	構造化プロパティ。
trails	number	ブーストまたはバグのコンポーネント・モデル数。
set_ensemble_method	Voting HighestProbability HighestMeanProbability	カテゴリー型対象のデフォルト結合ルール。
range_ensemble_method	Mean Median	連続型対象のデフォルト結合ルール。
large_boost	フラグ	特に大きなデータセットのブースティングを適用します。

表 122. chaidnode プロパティ (続き)

chaidnode プロパティ	値	プロパティの説明
split_alpha	number	分割の有意水準。
merge_alpha	number	結合の有意水準。
bonferroni_adjustment	フラグ	Bonferroni メソッドを使用して有意確率値を調整。
split_merged_categories	フラグ	マージしたカテゴリーの再分割を許可。
chi_square	Pearson LR	カイ 2 乗統計の計算に使用される方法 (Pearson または尤度比)
epsilon	number	期待されるセル度数の最小変化量。
max_iterations	number	収束のための最大反復回数。
set_random_seed	integer	
seed	number	
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	
maximum_number_of_models	integer	

coxregnode プロパティ



Cox 線型回帰ノードを使用すると、打ち切りレコードの存在下でイベントまでの時間のデータの生存モデルを構築します。このモデルは、入力変数の特定の値に対して特定の時間 (T) に対象のイベントが発生した確率を予測する生存関数を生成します。

例

```
node = stream.create("coxreg", "My node")
node.setPropertyValue("survival_time", "tenure")
node.setPropertyValue("method", "BackwardsStepwise")
# Expert tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("removal_criterion", "Conditional")
node.setPropertyValue("survival", True)
```

表 123. coxregnode プロパティ

coxregnode プロパティ	値	プロパティの説明
survival_time	フィールド	Cox 回帰モデルは 生存時間のある 1 つのフィールドを使用します。

表 123. coxregnode プロパティ (続き)

coxregnode プロパティ	値	プロパティの説明
target	フィールド	Cox 回帰モデルは 1 つの対象フィールドおよび 1 つ以上の入力フィールドを使用します。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
method	Enter Stepwise BackwardsStepwise	
groups	フィールド	
model_type	MainEffects Custom	
custom_terms	["BP*Sex" "BP*Age"]	
mode	Expert Simple	
max_iterations	number	
p_converge	1.0E-4 1.0E-5 1.0E-6 1.0E-7 1.0E-8 0	
p_converge	1.0E-4 1.0E-5 1.0E-6 1.0E-7 1.0E-8 0	

表 123. coxregnode プロパティ (続き)

coxregnode プロパティ	値	プロパティの説明
l_converge	1.0E-1 1.0E-2 1.0E-3 1.0E-4 1.0E-5 0	
removal_criterion	LR Wald Conditional	
probability_entry	<i>number</i>	
probability_removal	<i>number</i>	
output_display	EachStep LastStep	
ci_enable	フラグ	
ci_value	90 95 99	
correlation	フラグ	
display_baseline	フラグ	
survival	フラグ	
hazard	フラグ	
log_minus_log	フラグ	
one_minus_survival	フラグ	
separate_line	フィールド	
value	数値型 または 文字列	フィールドに対して値の指定がない場合、デフォルト・オプションの「Mean」をそのフィールドで使用します。

decisionlistnode プロパティ



ディシジョン・リスト・ノードは、母集団に関連する与えられた 2 値の結果の高いもしくは低い尤度を示すサブグループまたはセグメントを識別します。例えば、離れる可能性の少ないもしくはキャンペーンに好意的に答える可能性のある顧客を探すことができます。顧客区分を追加し、結果を比較するために他のモデルを並べて表示することによって、ビジネスに関する知識をモデルに導入することができます。ディシジョン・リスト・モデルは、ルール・リストから構成され、各ルールには条件と結果が含まれます。ルールは順番に適用され、一致する最初のルールで、結果が決まります。

例

```
node = stream.create("decisionlist", "My node")
node.setPropertyValue("search_direction", "Down")
node.setPropertyValue("target_value", 1)
node.setPropertyValue("max_rules", 4)
node.setPropertyValue("min_group_size_pct", 15)
```

表 124. decisionlistnode プロパティ

decisionlistnode プロパティ	値	プロパティの説明
target	フィールド	ディシジョン・リスト・モデルは 1 つの対象フィールドおよび 1 つ以上の入力フィールドを使用します。度数フィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
model_output_type	Model InteractiveBuilder	
search_direction	Up Down	セグメントの検索に関連します。Up は、高い確率の検索、Down は低い確率の検索と同じです。
target_value	string	指定しない場合は、フラグには真の値が想定されます。
max_rules	integer	残りを除外するセグメントの最大数
min_group_size	integer	最小セグメント・サイズ:
min_group_size_pct	number	最小セグメント・サイズ (パーセントとして)。
confidence_level	number	セグメント定義に追加するためにふさわしくするために、応答の尤度を向上するために入力フィールドが持つ最小しきい値。
max_segments_per_rule	integer	
mode	Simple Expert	

表 124. *decisionlistnode* プロパティ (続き)

decisionlistnode プロパティ	値	プロパティの説明
bin_method	EqualWidth EqualCount	
bin_count	number	
max_models_per_cycle	integer	リストの検索幅。
max_rules_per_cycle	integer	セグメント・ルールを検索幅。
segment_growth	number	
include_missing	フラグ	
final_results_only	フラグ	
reuse_fields	フラグ	属性 (ルールに表示される入力フィールド) の再使用を許可します。
max_alternatives	integer	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

discriminantnode プロパティ



判別分析では、ロジスティック回帰よりも厳密な仮定が行われますが、これらの仮定が満たされると、ロジスティック回帰分析に対する有益な代替手段または補足手段になる可能性があります。

例

```
node = stream.create("discriminant", "My node")
node.setPropertyValue("target", "custcat")
node.setPropertyValue("use_partitioned_data", False)
node.setPropertyValue("method", "Stepwise")
```

表 125. *discriminantnode* プロパティ

discriminantnode プロパティ	値	プロパティの説明
target	フィールド	判別分析 モデルは単一の対象フィールドおよび 1 つ以上の入力フィールドを使用します。重みフィールドおよび度数フィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。

表 125. *discriminantnode* プロパティ (続き)

discriminantnode プロパティ	値	プロパティの説明
method	Enter Stepwise	
mode	Simple Expert	
prior_probabilities	AllEqual ComputeFromSizes	
covariance_matrix	WithinGroups SeparateGroups	
means	フラグ	「詳細出力」ダイアログ・ボックスの統計オプション
univariate_anovas	フラグ	
box_m	フラグ	
within_group_covariance	フラグ	
within_groups_correlation	フラグ	
separate_groups_covariance	フラグ	
total_covariance	フラグ	
fishers	フラグ	
unstandardized	フラグ	
casewise_results	フラグ	「詳細出力」ダイアログ・ボックスの統計オプション
limit_to_first	number	デフォルト値は 10 です。
summary_table	フラグ	
leave_one_classification	フラグ	
combined_groups	フラグ	
separate_groups_covariance	フラグ	グループ別共分散行列オプション
territorial_map	フラグ	
combined_groups	フラグ	結合グループ散布図オプション
separate_groups	フラグ	グループ別散布図オプション
summary_of_steps	フラグ	
F_pairwise	フラグ	

表 125. *discriminantnode* プロパティ (続き)

discriminantnode プロパティ	値	プロパティの説明
stepwise_method	WilksLambda UnexplainedVariance MahalanobisDistance SmallestF RaosV	
V_to_enter	<i>number</i>	
criteria	UseValue UseProbability	
F_value_entry	<i>number</i>	デフォルト値は 3.84 です。
F_value_removal	<i>number</i>	デフォルト値は 2.71 です。
probability_entry	<i>number</i>	デフォルト値は 0.05 です。
probability_removal	<i>number</i>	デフォルト値は 0.10 です。
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

extensionmodelnode プロパティ



拡張モデル・ノードを使用すると、R スクリプトまたは Python for spark スクリプトを実行して、結果の作成およびスコアリングができます。

Python for Spark の例

```
##### script example for Python for Spark
import modeler.api
stream = modeler.script.stream()
node = stream.create("extension_build", "extension_build")
node.setPropertyValue("syntax_type", "Python")

build_script = """
import json
import spss.pyspark.runtime
```



```

from pyspark.mllib.regression import LabeledPoint
from pyspark.mllib.linalg import DenseVector
from pyspark.mllib.tree import DecisionTree

cxt = spss.pyspark.runtime.getContext()
df = cxt.getSparkInputData()
schema = df.dtypes[:]

target = "Drug"
predictors = ["Age", "BP", "Sex", "Cholesterol", "Na", "K"]

def metaMap(row,schema):
    col = 0
    meta = []
    for (cname, ctype) in schema:
        if ctype == 'string':
            meta.append(set([row[col]]))
        else:
            meta.append((row[col],row[col]))
        col += 1
    return meta

def metaReduce(meta1,meta2,schema):
    col = 0
    meta = []
    for (cname, ctype) in schema:
        if ctype == 'string':
            meta.append(meta1[col].union(meta2[col]))
        else:
            meta.append((min(meta1[col][0],meta2[col][0]),max(meta1[col][1],meta2[col][1])))
        col += 1
    return meta

metadata = df.rdd.map(lambda row: metaMap(row,schema)).reduce(lambda x,y:metaReduce(x,y,schema))

def setToList(v):
    if isinstance(v,set):
        return list(v)
    return v

metadata = map(lambda x: setToList(x), metadata)
print metadata

lookup = {}
for i in range(0,len(schema)):
    lookup[schema[i][0]] = i

def row2LabeledPoint(dm,lookup,target,predictors,row):
    target_index = lookup[target]
    tval = dm[target_index].index(row[target_index])
    pvals = []
    for predictor in predictors:
        predictor_index = lookup[predictor]
        if isinstance(dm[predictor_index],list):
            pval = dm[predictor_index].index(row[predictor_index])
        else:
            pval = row[predictor_index]
        pvals.append(pval)
    return LabeledPoint(tval,DenseVector(pvals))

# count number of target classes
predictorClassCount = len(metadata[lookup[target]])

# define function to extract categorical predictor information from datamodel
def getCategoricalFeatureInfo(dm,lookup,predictors):
    info = {}
    for i in range(0,len(predictors)):
        predictor = predictors[i]
        predictor_index = lookup[predictor]
        if isinstance(dm[predictor_index],list):
            info[i] = len(dm[predictor_index])
    return info

# convert dataframe to an RDD containing LabeledPoint
lps = df.rdd.map(lambda row: row2LabeledPoint(metadata,lookup,target,predictors,row))

treeModel = DecisionTree.trainClassifier(
    lps,
    numClasses=predictorClassCount,
    categoricalFeaturesInfo=getCategoricalFeatureInfo(metadata, lookup, predictors),
    impurity='gini',
    maxDepth=5,

```

```

maxBins=100)

_outputPath = cxt.createTemporaryFolder()
treeModel.save(cxt.getSparkContext(), _outputPath)
cxt.setModelContentFromPath("TreeModel", _outputPath)
cxt.setModelContentFromString("model.dm", json.dumps(metadata), mimeType="application/json")\
    .setModelContentFromString("model.structure", treeModel.toDebugString())

"""

node.setPropertyValue("python_build_syntax", build_script)

```

R の例

```

#### script example for R
node.setPropertyValue("syntax_type", "R")
node.setPropertyValue("r_build_syntax", """modelerModel <-
lm(modelerData$Na~modelerData$K,modelerData)
modelerDataModel
modelerModel
""")

```

表 126. *extensionmodelnode* プロパティ

extensionmodelnode プロパティ	値	プロパティの説明
syntax_type	R Python	R または Python のどちらのスクリプトを実行するか指定します (R がデフォルトです)。
r_build_syntax	string	モデル作成用の R スクリプト・シンタックス。
r_score_syntax	string	モデル・スコアリング用の R スクリプト・シンタックス。
python_build_syntax	string	モデル作成用の Python スクリプト・シンタックス。
python_score_syntax	string	モデル・スコアリング用の Python スクリプト・シンタックス。
convert_flags	StringsAndDoubles LogicalValues	フラグ型フィールドを変換するためのオプション。
convert_missing	フラグ	欠損値を R NA 値に変換するオプションです。
convert_datetime	フラグ	日付形式または日付/時刻形式の変数を R の日付/時刻形式に変換するためのオプション。
convert_datetime_class	POSIXct POSIXlt	日付形式または日付/時刻形式の変数のうち、どの形式の変数を変換するかを指定するためのオプション。
output_html	フラグ	R モデル・ナゲットのタブにグラフを表示するためのオプション。
output_text	フラグ	R モデル・ナゲットのタブに R コンソールのテキスト出力を書き込むためのオプション。

factornode プロパティ



PCA/因子ノードには、データの複雑性を整理する強力なデータ分解手法が用意されています。主成分分析 (PCA) : 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 (直角に交わる) 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。因子分析 : 一連の観測フィールド内の相関パターンを説明する基本因子が識別されます。どちらの手法でも、元のフィールド・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の検出が目標です。

例

```
node = stream.create("factor", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("inputs", ["BP", "Na", "K"])
node.setPropertyValue("partition", "Test")
# "Model" tab
node.setPropertyValue("use_model_name", True)
node.setPropertyValue("model_name", "Factor_Age")
node.setPropertyValue("use_partitioned_data", False)
node.setPropertyValue("method", "GLS")
# Expert options
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("complete_records", True)
node.setPropertyValue("matrix", "Covariance")
node.setPropertyValue("max_iterations", 30)
node.setPropertyValue("extract_factors", "ByFactors")
node.setPropertyValue("min_eigenvalue", 3.0)
node.setPropertyValue("max_factor", 7)
node.setPropertyValue("sort_values", True)
node.setPropertyValue("hide_values", True)
node.setPropertyValue("hide_below", 0.7)
# "Rotation" section
node.setPropertyValue("rotation", "DirectOblimin")
node.setPropertyValue("delta", 0.3)
node.setPropertyValue("kappa", 7.0)
```

表 127. factornode プロパティ		
factornode プロパティ	値	プロパティの説明
inputs	[field1 ... fieldN]	主成分分析/因子モデルは対象フィールドでなく、入力フィールドのリストを使用します。重みフィールドおよび度数フィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。

表 127. *factornode* プロパティ (続き)

factornode プロパティ	値	プロパティの説明
method	PC ULS GLS ML PAF Alpha Image	
mode	Simple Expert	
max_iterations	<i>number</i>	
complete_records	フラグ	
matrix	Correlation Covariance	
extract_factors	ByEigenvalues ByFactors	
min_eigenvalue	<i>number</i>	
max_factor	<i>number</i>	
rotation	None Varimax DirectOblimin Equamax Quartimax Promax	
delta	<i>number</i>	rotation で DirectOblimin を選択した場合、delta の値を指定できる。 値を指定しない場合は、delta のデフォルト値を使用。

表 127. *factornode* プロパティ (続き)

factornode プロパティ	値	プロパティの説明
kappa	<i>number</i>	rotation で Promax を選択した場合、kappa の値を指定できる。 値を指定しない場合は、kappa のデフォルト値を使用。
sort_values	フラグ	
hide_values	フラグ	
hide_below	<i>number</i>	

featureselectionnode プロパティ



特徴量選択ノードで、(欠損値の割合などの) 諸基準に基づいて入力フィールドをスクリーニングして削除にかけ、指定した目標に相対的な残りの入力フィールドの重要度をランク付けします。例えば、数百の潜在的入力フィールドを含むデータセットがあるとして、患者予後のモデリングにはどれが役に立つのでしょうか？

例

```
node = stream.create("featureselection", "My node")
node.setPropertyValue("screen_single_category", True)
node.setPropertyValue("max_single_category", 95)
node.setPropertyValue("screen_missing_values", True)
node.setPropertyValue("max_missing_values", 80)
node.setPropertyValue("criteria", "Likelihood")
node.setPropertyValue("unimportant_below", 0.8)
node.setPropertyValue("important_above", 0.9)
node.setPropertyValue("important_label", "Check Me Out!")
node.setPropertyValue("selection_mode", "TopN")
node.setPropertyValue("top_n", 15)
```

特徴量選択モデルの作成および適用の詳細な例は、4 ページの『スタンドアロン スクリプトの例:特徴量選択モデルの生成』を参照してください。

表 128. *featureselectionnode* プロパティ

featureselectionnode プロパティ	値	プロパティの説明
target	フィールド	特徴量選択モデルは指定対象に関連した予測フィールドをランク付けします。重みフィールドおよび度数フィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
screen_single_category	フラグ	True の場合、総レコード数に比べ同じカテゴリーに多くかたよったレコードを持つフィールドを選別します。
max_single_category	<i>number</i>	screen_single_category が True の場合に使用される閾値を指定します。

表 128. *featureselectionnode* プロパティ (続き)

featureselectionnode プロパティ	値	プロパティの説明
screen_missing_values	フラグ	True の場合、レコードの総数のパーセントで表すレコード数になるまで、多すぎる欠損値フィールドをスクリーニング (選別) します。
max_missing_values	<i>number</i>	
screen_num_categories	フラグ	True の場合、レコードの総数に対して多すぎるカテゴリーを減らす目的で、フィールドをスクリーニング (選別) します。
max_num_categories	<i>number</i>	
screen_std_dev	フラグ	True の場合、指定された最小値以下の標準偏差で、フィールドをスクリーニング (選別) します。
min_std_dev	<i>number</i>	
screen_coeff_of_var	フラグ	True の場合、指定された最小値以下の分散係数で、フィールドをスクリーニング (選別) します。
min_coeff_of_var	<i>number</i>	
criteria	Pearson Likelihood CramersV Lambda	カテゴリー対象に対するカテゴリー予測値のランク付けのときに、重要な値が基準とする測定単位を指定します。
unimportant_below	<i>number</i>	重要、境界、非重要として変数をランク付けするときに使用される閾値 p を指定します。0.0 から 1.0 の値を指定します。
important_above	<i>number</i>	0.0 から 1.0 の値を指定します。
unimportant_label	<i>string</i>	非重要ランクのラベルを指定します。
marginal_label	<i>string</i>	
important_label	<i>string</i>	
selection_mode	ImportanceLevel ImportanceValue TopN	
select_important	フラグ	<code>selection_mode</code> が <code>ImportanceLevel</code> に設定されているときに、重要なフィールドを選択するかどうかを指定します。

表 128. *featureselectionnode* プロパティ (続き)

featureselectionnode プロパティ	値	プロパティの説明
select_marginal	フラグ	selection_mode が ImportanceLevel に設定されているときに、境界フィールドを選択するかどうかを指定します。
select_unimportant	フラグ	selection_mode が ImportanceLevel に設定されているときに、重要でないフィールドを選択するかどうかを指定します。
importance_value	number	selection_mode が ImportanceValue に設定されているときに、使用する分割値を指定します。0 から 100 の値を指定します。
top_n	integer	selection_mode が TopN に設定されているときに、使用する分割値を指定します。0 から 1000 の値を指定します。

genlinnode プロパティ



一般化線型モデルは、指定したリンク関数によって従属変数が因子および共変量と線型関係になるよう、一般線型モデルを拡張したものです。さらに、このモデルは、非正規分布の従属変数に対応します。線型回帰、ロジスティック回帰、度数データの対数線型モデル、区間打ち切り生存モデルなど、多数の統計モデルの機能をカバーします。

例

```
node = stream.create("genlin", "My node")
node.setPropertyValue("model_type", "MainAndAllTwoWayEffects")
node.setPropertyValue("offset_type", "Variable")
node.setPropertyValue("offset_field", "Claimant")
```

表 129. *genlinnode* プロパティ

genlinnode プロパティ	値	プロパティの説明
target	フィールド	一般化線型モデルは、名義型またはフラグ型の 1 つの対象フィールドおよび 1 つ以上の入力フィールドが必要です。重みフィールドも指定できます。詳しくは、トピック 217 ページ の『 一般的なモデル作成ノードのプロパティ 』を参照してください。
use_weight	フラグ	
weight_field	フィールド	フィールドのデータ型は連続型だけです。
target_represents_trials	フラグ	

表 129. genlinnode プロパティ (続き)

genlinnode プロパティ	値	プロパティの説明
trials_type	Variable FixedValue	
trials_field	フィールド	フィールドのデータ型はフラグ型または順序型です。
trials_number	number	デフォルト値は 10 です。
model_type	MainEffects MainAndAllTwoWayEffects	
offset_type	Variable FixedValue	
offset_field	フィールド	フィールドのデータ型は連続型だけです。
offset_value	number	実数である必要があります。
base_category	Last First	
include_intercept	フラグ	
mode	Simple Expert	
distribution	BINOMIAL GAMMA IGAUSS NEGBIN NORMAL POISSON TWEEDIE MULTINOMIAL	IGAUSS: 逆ガウス。 NEGBIN: 負の 2 項分布。
negbin_para_type	Specify Estimate	

表 129. genlinnode プロパティ (続き)

genlinnode プロパティ	値	プロパティの説明
negbin_parameter	<i>number</i>	デフォルト値は 1 です。負でない実数が含まれている必要があります。
tweedie_parameter	<i>number</i>	
link_function	IDENTITY CLOGLOG LOG LOGC LOGIT NEGBIN NLOGLOG ODDSPower PROBIT POWER CUMCAUCHIT CUMCLOGLOG CUMLOGIT CUMNLOGLOG CUMPROBIT	CLOGLOG: 補ログ・マイナス・ログ。 LOGC: 補対数。 NEGBIN: 負の 2 項分布。 NLOGLOG: 負ログ・マイナス・ログ。 CUMCAUCHIT: 累積コーチット。 CUMCLOGLOG: 累積補ログ・マイナス・ログ。 CUMLOGIT: 累積ロジット。 CUMNLOGLOG: 累積負ログ・マイナス・ログ。 CUMPROBIT: 累積プロビット。
power	<i>number</i>	値は 0 でない実数である必要があります。
method	Hybrid Fisher NewtonRaphson	
max_fisher_iterations	<i>number</i>	デフォルト値は 1 です。正の整数値だけが使用できます。

表 129. genlinnode プロパティ (続き)

genlinnode プロパティ	値	プロパティの説明
scale_method	MaxLikelihoodEstimate Deviance PearsonChiSquare FixedValue	
scale_value	number	デフォルト値は 1 です。0 を超える必要があります。
covariance_matrix	ModelEstimator RobustEstimator	
max_iterations	number	デフォルト値は 100 です。0 以上の整数だけを使用できます。
max_step_halving	number	デフォルト値は 5 です。正の整数値だけが使用できます。
check_separation	フラグ	
start_iteration	number	デフォルト値は 20 です。正の整数値だけが使用できます。
estimates_change	フラグ	
estimates_change_min	number	デフォルト値は 1E-006 です。正の数値だけが使用できます。
estimates_change_type	Absolute Relative	
loglikelihood_change	フラグ	
loglikelihood_change_min	number	正の数値だけが使用できます。
loglikelihood_change_type	Absolute Relative	
hessian_convergence	フラグ	
hessian_convergence_min	number	正の数値だけが使用できます。
hessian_convergence_type	Absolute Relative	
case_summary	フラグ	
contrast_matrices	フラグ	
descriptive_statistics	フラグ	
estimable_functions	フラグ	
model_info	フラグ	

表 129. *genlinnode* プロパティ (続き)

genlinnode プロパティ	値	プロパティの説明
iteration_history	フラグ	
goodness_of_fit	フラグ	
print_interval	<i>number</i>	デフォルト値は 1 です。正の整数である必要があります。
model_summary	フラグ	
lagrange_multiplier	フラグ	
parameter_estimates	フラグ	
include_exponential	フラグ	
covariance_estimates	フラグ	
correlation_estimates	フラグ	
analysis_type	TypeI TypeIII TypeIAndTypeIII	
statistics	Wald LR	
citype	Wald Profile	
tolerancelevel	<i>number</i>	デフォルト値は 0.0001 です。
confidence_interval	<i>number</i>	デフォルト値は 95 です。
loglikelihood_function	Full Kernel	
singularity_tolerance	1E-007 1E-008 1E-009 1E-010 1E-011 1E-012	

表 129. genlinnode プロパティ (続き)

genlinnode プロパティ	値	プロパティの説明
value_order	Ascending Descending DataOrder	
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

glmmnode プロパティ



一般化線型混合モデル (GLMN) は線型モデルを拡張したため、対象が非正規分布となる場合があります。指定されたリンク関数を介して因子および共変量に線形に関連し、観測が相関できるようになりました。一般化線型混合モデルには、単純な線型回帰から、正規分布していない縦断的データを取り扱う複雑なマルチレベル・モデルまで、さまざまなモデルが含まれます。

表 130. glmmnode プロパティ

glmmnode プロパティ	値	プロパティの説明
residual_subject_spec	構造化	指定したカテゴリー型フィールドの組み合わせにより、データセット内の被験者が一意に定義されることが必要です。
repeated_measures	構造化	反復する観察の特定に使用されるフィールド。
residual_group_spec	[field1 ... fieldN]	反復効果共変量パラメーターの独立セットを定義するフィールド。

表 130. *glmmnode* プロパティ (続き)

glmmnode プロパティ	値	プロパティの説明
residual_covariance_type	Diagonal AR1 ARMA11 COMPOUND_SYMMETRY IDENTITY TOEPLITZ UNSTRUCTURED VARIANCE_COMPONENTS	残差の共変量構造を指定します。
custom_target	フラグ	上流のノードで定義された対象を使用するか (<code>false</code>) または <code>target_field</code> によって指定されたカスタム対象を使用するか (<code>true</code>) を定義します。
target_field	フィールド	<code>custom_target</code> が <code>true</code> の場合対象として使用するフィールド。
use_trials	フラグ	試行回数を指定する追加フィールド又は値を、対象フィールドが一連の試行が発生する様々なイベントである場合に使用するかどうかを示します。デフォルトは <code>false</code> です。
use_field_or_value	Field Value	フィールドまたは値を使用して試行回数を指定するかどうかを示します。
trials_field	フィールド	試行回数の指定に使用するフィールド。
trials_value	<i>integer</i>	試行回数の指定に使用する値。指定する場合、最小値は 1 です。
use_custom_target_reference	フラグ	カスタム参照カテゴリーをカテゴリー型対象に使用するかどうかを示します。デフォルトは <code>false</code> です。
target_reference_value	<i>string</i>	<code>use_custom_target_reference</code> が <code>true</code> の場合使用する参照カテゴリー。

表 130. *glmmnode* プロパティ (続き)

glmmnode プロパティ	値	プロパティの説明
dist_link_combination	Nominal Logit GammaLog BinomialLogit PoissonLog BinomialProbit NegbinLog BinomialLogC Custom	対象の値の分布に関する一般モデル。 Custom を選択して、 target_distribution で提供された リストから分布を指定します。
target_distribution	Normal Binomial Multinomial Gamma Inverse NegativeBinomial Poisson	dist_link_combination が Custom の場合の対象の値の分布。

表 130. *glmmnode* プロパティ (続き)

glmmnode プロパティ	値	プロパティの説明
link_function_type	Identity LogC Log CLOGLOG Logit NLOGLOG PROBIT POWER CAUCHIT	ターゲットに関連付けるリンク関数値を予測値に変換します。 target_distribution の場合 Binomial 以下のいずれかを使用できます。 リストされたリンク関数のうちの 1 つ。 target_distribution の場合 Multinomial 以下を使用できます。 CLOGLOG, CAUCHIT, LOGIT, NLOGLOG、または PROBIT。 target_distribution の場合 Binomial 以外のもの、または Multinomial 以下を使用できます。 IDENTITY、LOG、または POWER。
link_function_param	number	使用するリンク関数パラメーター値。 normal_link_function または link_function_type が POWER の場合のみ適用されます。
use_predefined_inputs	フラグ	固定効果フィールドを、入力フィールド (true) として上流に定義されたフィールドにするか、fixed_effects_list (false) から定義されたフィールドにするかを示します。デフォルトは false です。
fixed_effects_list	構造化	use_predefined_inputs が false の場合、固定効果フィールドとして使用する入力フィールドを指定します。
use_intercept	フラグ	true (デフォルト) の場合、モデルに定数項を含みます。
random_effects_list	構造化	ランダム効果として指定するフィールドのリスト。
regression_weight_field	フィールド	分析の重みフィールドとして使用するフィールド。
use_offset	None offset_value offset_field	オフセットを指定する方法を示します。値 None は、オフセットが使用されないことを意味します。
offset_value	number	use_offset が offset_value の場合オフセットに使用する値。

表 130. *glmmnode* プロパティ (続き)

glmmnode プロパティ	値	プロパティの説明
offset_field	フィールド	use_offset が offset_field の場合オフセット値に使用する値。
target_category_order	Ascending Descending Data	カテゴリー型対象のソート順。値 Data は、データ内のソート順を使用するよう指定します。デフォルトは Ascending です。
inputs_category_order	Ascending Descending Data	Sorting order for categorical predictors. 値 Data は、データ内のソート順を使用するよう指定します。デフォルトは Ascending です。
max_iterations	integer	アルゴリズムで実行される反復の最大回数です。負の数ではない整数。デフォルト値は 100 です。
confidence_level	integer	モデル係数の区間推定の計算に使用する確信度。負の数ではない整数。最小値は 100、デフォルト値は 95 です。
degrees_of_freedom_method	Fixed Varied	自由度が有意性検定に計算される方法を指定します。
test_fixed_effects_coefficients	Model Robust	パラメーター推定共変量マトリックスを計算する方法。
use_p_converge	フラグ	パラメーター収束のオプション。
p_converge	number	空白または任意の正の値。
p_converge_type	Absolute Relative	
use_l_converge	フラグ	対数尤度収束のオプション。
l_converge	number	空白または任意の正の値。
l_converge_type	Absolute Relative	
use_h_converge	フラグ	Hessian 収束のオプション。
h_converge	number	空白または任意の正の値。
h_converge_type	Absolute Relative	
max_fisher_steps	integer	
singularity_tolerance	number	

表 130. glmmnode プロパティ (続き)

glmmnode プロパティ	値	プロパティの説明
use_model_name	フラグ	モデルのカスタム名を指定するか (true)、システム生成名を使用するか (false) を示します。デフォルトは false です。
model_name	string	use_model_name が true のときに、使用するモデルを指定します。
confidence	onProbability onIncrease	スコアリングの確信度を計算する基準 (最も高い予測確率、または最も高い予測確率と 2 番目に高い予測確率との差)。
score_category_probabilities	フラグ	true の場合、カテゴリー型対象の予測確率を生成します。デフォルトは false です。
max_categories	integer	score_category_probabilities が true のときに、使用するカテゴリーの最大数を指定します。
score_propensity	フラグ	true の場合、フィールドの「true」の結果の確率を示すフラグ型対象フィールドの傾向スコアを生成します。
emeans	Structure	固定効果リストの各カテゴリー型フィールドについて、推定周辺平均を生成するかどうかを指定します。
covariance_list	Structure	固定効果リストの各カテゴリー型フィールドについて、推定周辺平均を計算する場合に平均値を使用するかカスタム値を使用するかを指定します。
mean_scale	Original Transformed	対象の元の尺度に基づいて (デフォルト)、またはリンク関数変換に基づいて推定周辺平均を計算するかどうかを指定します。
comparison_adjustment_method	LSD SEQBONFERRONI SEQSIDAK	複数の対比で仮定検定を実行する場合に使用する調整方法。

gle プロパティ



GLE は、対象を非正規分布とできるように線型モデルを拡張したものであり、指定されたリンク関数を介して因子および共変量に線形に関連し、観測が相関できるようにしました。一般化線型混合モデルには、単純な線型回帰から、正規分布していない縦断的データを取り扱う複雑なマルチレベル・モデルまで、さまざまなモデルが含まれます。

表 131. gle プロパティ

gle プロパティ	値	プロパティの説明
custom_target	フラグ	上流のノードで定義された対象を使用するか (false) または target_field によって指定されたカスタム対象を使用するか (true) を定義します。
target_field	フィールド	custom_target が true の場合対象として使用するフィールド。
use_trials	フラグ	試行回数を指定する追加フィールド又は値を、対象フィールドが一連の試行が発生する様々なイベントである場合に使用するかどうかを示します。デフォルトは false です。
use_trials_field_or_value	Field Value	フィールドまたは値を使用して試行回数を指定するかどうかを示します。
trials_field	フィールド	試行回数の指定に使用するフィールド。
trials_value	integer	試行回数の指定に使用する値。指定する場合、最小値は 1 です。
use_custom_target_reference	フラグ	カスタム参照カテゴリをカテゴリ型対象に使用するかどうかを示します。デフォルトは false です。
target_reference_value	string	use_custom_target_reference が true の場合使用する参照カテゴリ。
dist_link_combination	NormalIdentity GammaLog PoissonLog NegbinLog TweedieIdentity NominalLogit BinomialLogit BinomialProbit BinomialLogC CUSTOM	対象の値の分布に関する一般モデル。 target_distribution で提供されるリストから分布を選択するには CUSTOM を選択します。

表 131. gle プロパティ (続き)

gle プロパティ	値	プロパティの説明
target_distribution	Normal Binomial Multinomial Gamma INVERSE_GAUSS NEG_BINOMIAL Poisson TWEEDIE UNKNOWN	dist_link_combination が Custom の場合の対象の値の分布。

表 131. gle プロパティ (続き)

gle プロパティ	値	プロパティの説明
link_function_type	UNKNOWN	対象値を予測値に関連付けるリンク関数。 target_distribution が Binomial の場合、以下を使用できます。
	IDENTITY	
	LOG	
	LOGIT	IDENTITY
	PROBIT	LOG
	COMPL_LOG_LOG	LOGIT
	POWER	PROBIT
	LOG_COMPL	COMPL_LOG_LOG
	NEG_LOG_LOG	POWER
	ODDS_POWER	LOG_COMPL
	NEG_BINOMIAL	NEG_LOG_LOG
	GEN_LOGIT	ODDS_POWER
	CUMUL_LOGIT	target_distribution が NEG_BINOMIAL の場合、以下を使用できます。
	CUMUL_PROBIT	
	CUMUL_COMPL_LOG_LOG	NEG_BINOMIAL.
	CUMUL_NEG_LOG_LOG	target_distribution が UNKNOWN の場合、以下を使用できます。
	CUMUL_CAUCHIT	
		GEN_LOGIT
		CUMUL_LOGIT
		CUMUL_PROBIT
		CUMUL_COMPL_LOG_LOG
		CUMUL_NEG_LOG_LOG
		CUMUL_CAUCHIT

表 131. gle プロパティ (続き)

gle プロパティ	値	プロパティの説明
link_function_param	<i>number</i>	使用する Tweedie パラメータ。 normal_link_function または link_function_type が POWER の場合の み適用されます。
tweedie_param	<i>number</i>	使用するリンク関数パラメーター値。 dist_link_combination が TweedieIdentity に設定されているか、ま たは link_function_type が TWEEDIE の 場合にのみ適用できます。
use_predefined_inputs	フラグ	モデル効果フィールドを入力フィールドとし て上流で定義されたフィールドとするか (true) fixed_effects_list のフィールド とするか (false) を指定します。
model_effects_list	構造化	use_predefined_inputs が false の場 合、モデル効果フィールドとして使用する入 力フィールドを指定します。
use_intercept	フラグ	true (デフォルト) の場合、モデルに定数項を 含みます。
regression_weight_field	フィールド	分析の重みフィールドとして使用するフィー ルド。
use_offset	None Value Variable	オフセットを指定する方法を示します。 値 None は、オフセットが使用されないことを意 味します。
offset_value	<i>number</i>	use_offset が offset_value の場合オフ セットに使用する値。
offset_field	フィールド	use_offset が offset_field の場合オフ セット値に使用する値。
target_category_order	Ascending Descending	カテゴリー型対象のソート順。 デフォルトは Ascending です。
inputs_category_order	Ascending Descending	Sorting order for categorical predictors. デフ ォルトは Ascending です。
max_iterations	<i>integer</i>	アルゴリズムで実行される反復の最大回数で す。 負の数ではない整数。デフォルト値は 100 です。
confidence_level	<i>number</i>	モデル係数の区間推定の計算に使用する確信 度。 負の数ではない整数。最小値は 100、デ フォルト値は 95 です。
test_fixed_effects_coef fecients	Model Robust	パラメーター推定共変量マトリックスを計算 する方法。

表 131. gle プロパティ (続き)

gle プロパティ	値	プロパティの説明
detect_outliers	フラグ	true の場合、アルゴリズムで、多項分布を除くすべての分布に対する影響がある外れ値を検出します。
conduct_trend_analysis	フラグ	true の場合、アルゴリズムで散布図のトレンド分析を実行します。
estimation_method	FISHER_SCORING NEWTON_RAPHSON HYBRID	最尤法推定アルゴリズムを指定します。
max_fisher_iterations	integer	FISHER_SCORING estimation_method を使用している場合の、最大反復回数。最小値は 0、最大値は 20 です。
scale_parameter_method	MLE FIXED DEVIANCE PEARSON_CHISQUARE	スケールパラメータの推定に使用する方法を指定します。
scale_value	number	scale_parameter_method が Fixed に設定されている場合にのみ使用できます。
negative_binomial_method	MLE FIXED	負の二項分布補助パラメータの推定の使用方法を指定します。
negative_binomial_value	number	negative_binomial_method が Fixed に設定されている場合にのみ使用できます。
non_neg_least_squares	フラグ	非負拘束付最小 2 乗法を実行するかどうか。デフォルトは false です。
use_p_converge	フラグ	パラメーター収束のオプション。
p_converge	number	空白または任意の正の値。
p_converge_type	フラグ	True = 絶対値、False = 相対値
use_l_converge	フラグ	対数尤度収束のオプション。
l_converge	number	空白または任意の正の値。
l_converge_type	フラグ	True = 絶対値、False = 相対値
use_h_converge	フラグ	Hessian 収束のオプション。
h_converge	number	空白または任意の正の値。
h_converge_type	フラグ	True = 絶対値、False = 相対値
max_iterations	integer	アルゴリズムで実行される反復の最大回数です。負の数ではない整数。デフォルト値は 100 です。
sing_tolerance	integer	
use_model_selection	フラグ	パラメータしきい値とモデルの選択方法コントロールを有効にします。

表 131. gle プロパティ (続き)

gle プロパティ	値	プロパティの説明
method	LASSO ELASTIC_NET FORWARD_STEPWISE RIDGE	モデルの選択方法、または Ridge を使用している場合は正規化方法を決定します。
detect_two_way_interactions	フラグ	True の場合、モデルにより入力フィールド間の 2 要因の交互作用が自動的に検出されます。 このコントロールは、モデルが主効果のみ (ユーザーが高次元効果を作成していない) であり、かつ選択された method が Forward Stepwise、Lasso、または Elastic Net の場合にのみ有効にしてください。
automatic_penalty_params	フラグ	モデル選択の method が Lasso または Elastic Net の場合のみ使用可能です。 この機能を使用して、Lasso または Elastic Net 変数選択方法に関連付けられたペナルティ パラメータを入力します。 True の場合、デフォルト値が使用されます。False の場合、ペナルティ パラメータが有効になり、カスタム値を入力できます。
lasso_penalty_param	number	モデル選択の method が Lasso または Elastic Net であり、automatic_penalty_params が False の場合にのみ使用できます。Lasso のペナルティ パラメータを指定します。
elastic_net_penalty_param1	number	モデル選択の method が Lasso または Elastic Net であり、automatic_penalty_params が False の場合にのみ使用できます。Elastic Net パラメータ 1 のペナルティ パラメータを指定します。
elastic_net_penalty_param2	number	モデル選択の method が Lasso または Elastic Net であり、automatic_penalty_params が False の場合にのみ使用できます。Elastic Net パラメータ 2 のペナルティ パラメータを指定します。
probability_entry	number	選択された method が Forward Stepwise の場合にのみ使用できます。効果の包含に関する F 統計基準の有意水準を指定します。
probability_removal	number	選択された method が Forward Stepwise の場合にのみ使用できます。効果の除去に関する F 統計基準の有意水準を指定します。

表 131. gle プロパティ (続き)

gle プロパティ	値	プロパティの説明
use_max_effects	フラグ	選択された method が Forward Stepwise の場合にのみ使用できます。 max_effects コントロールを有効にします。 False の場合、包含する効果のデフォルト数が、モデルに提供される効果の総数から切片を引いたものと等しくなければなりません。
max_effects	integer	変数増加ステップワイズ法作成方法を使用する場合の効果の最大数を指定します。
use_max_steps	フラグ	max_steps コントロールを有効にします。 False の場合、ステップのデフォルト数は、モデルに提供された効果の数から切片を除外したものの 3 倍と等しくなければなりません。
max_steps	integer	変数増加ステップワイズ法作成 method を使用する際に取るステップの最大数を指定します。
use_model_name	フラグ	モデルのカスタム名を指定するか (true)、システム生成名を使用するか (false) を示します。デフォルトは false です。
model_name	string	use_model_name が true のときに、使用するモデルを指定します。
usePI	フラグ	true の場合、予測変数の重要度が計算されます。

kmeansnode プロパティ



K-Means ノードで、データ・セットが異なるグループ (つまりクラスター) へ、クラスティングされます。この方法で、固定数のクラスターを定義し、クラスターにレコードを繰り返し割り当てて、これ以上調整してもモデルが改善されなくなるまで、クラスターの中心を調整します。K-means では、結果を予測するのではなく、入力フィールドのセット内のパターンを明らかにするために、「教師なし学習」として知られるプロセスが使用されます。

例

```
node = stream.create("kmeans", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("inputs", ["Cholesterol", "BP", "Drug", "Na", "K",
"Age"])
# "Model" tab
node.setPropertyValue("use_model_name", True)
node.setPropertyValue("model_name", "Kmeans_allinputs")
node.setPropertyValue("num_clusters", 9)
node.setPropertyValue("gen_distance", True)
node.setPropertyValue("cluster_label", "Number")
```



```

node.setPropertyValue("label_prefix", "Kmeans_")
node.setPropertyValue("optimize", "Speed")
# "Expert" tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("stop_on", "Custom")
node.setPropertyValue("max_iterations", 10)
node.setPropertyValue("tolerance", 3.0)
node.setPropertyValue("encoding_value", 0.3)

```

表 132. *kmeansnode* プロパティ

kmeansnode プロパティ	値	プロパティの説明
inputs	[field1 ... fieldN]	K-means モデルは入力フィールドのセットでクラスター分析を行います。対象フィールドは使用しません。重みフィールドおよび度数フィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
num_clusters	number	
gen_distance	フラグ	
cluster_label	String Number	
label_prefix	string	
mode	Simple Expert	
stop_on	Default Custom	
max_iterations	number	
tolerance	number	
encoding_value	number	
optimize	Speed Memory	モデル作成が速度とメモリーのどちらにより最適化されるかを指定します。

kmeansasnode プロパティ



K-Means は、最も一般的に使用されるクラスタリングアルゴリズムの 1 つです。このアルゴリズムは、データポイントをクラスタリングして、事前定義された数のクラスターを作成します。スポンサー モデラー の K-Means-AS ノードは Spark で実装されています。K-Means アルゴリズムについて詳しくは、<https://spark.apache.org/docs/2.2.0/ml-clustering.html> を参照してください。K-Means-AS ノードでは、カテゴリ変数の場合にワンホットエンコーディングが自動的に実行されることに留意してください。

表 133. kmeansasnode プロパティ

kmeansasnode プロパティ	値	プロパティの説明
roleUse	string	定義済みの役割を使用するには predefined を指定し、ユーザー設定フィールドの割り当てを使用するには custom を指定します。デフォルトは predefined です。
autoModel	Boolean	新しく生成されたスコアリングフィールドにデフォルト名 (\$S-prediction) を使用するには true を指定し、カスタム名を使用するには false を指定します。デフォルトは true です。
features	フィールド	roleUse プロパティが custom に設定されている場合の入力用のフィールド名のリスト。
name	string	autoModel プロパティが false に設定されている場合の、新しく生成されたスコアリングフィールドの名前。
clustersNum	integer	作成するクラスターの数。デフォルトは 5 です。
initMode	string	初期化アルゴリズム。指定できる値は、k-means または random です。デフォルトは k-means です。
initSteps	integer	initMode が k-means に設定されている場合の初期化ステップの数。デフォルトは 2 です。
advancedSettings	Boolean	次の 4 つのプロパティを使用可能にするには true を指定します。デフォルトは false です。
maxIteration	integer	クラスタリングの最大反復数。デフォルトは 20 です。
tolerance	string	反復を停止する許容度。可能な設定は、1.0E-1、1.0E-2、...、1.0E-6 です。デフォルトは 1.0E-4 です。
setSeed	Boolean	カスタム ランダム シードを使用するには true を指定します。デフォルトは false です。
randomSeed	integer	setSeed プロパティが true の場合のカスタム ランダム シード。

knnode プロパティ



k が整数である場合、 k 最近傍 (KNN) ノードは、新しいケースを、予測領域の新しいケースに最も近い k 個のオブジェクトのカテゴリまたは値と関連付けます。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。

例

```
node = stream.create("knn", "My node")
# Objectives tab
node.setPropertyValue("objective", "Custom")
# Settings tab - Neighbors panel
node.setPropertyValue("automatic_k_selection", False)
node.setPropertyValue("fixed_k", 2)
node.setPropertyValue("weight_by_importance", True)
# Settings tab - Analyze panel
node.setPropertyValue("save_distances", True)
```

表 134. knnnode プロパティ

knnnode プロパティ	値	プロパティの説明
analysis	PredictTarget	
	IdentifyNeighbors	
objective	Balance	
	Speed	
	Accuracy	
	Custom	
normalize_ranges	フラグ	
use_case_labels	フラグ	次のオプションを有効化するチェック・ボックス。
case_labels_field	フィールド	
identify_focal_cases	フラグ	次のオプションを有効化するチェック・ボックス。
focal_cases_field	フィールド	
automatic_k_selection	フラグ	
fixed_k	integer	automatic_k_selection が False の場合にのみ有効です。
minimum_k	integer	automatic_k_selection が True の場合にのみ有効です。
maximum_k	integer	
distance_computation	Euclidean	
	CityBlock	
weight_by_importance	フラグ	
range_predictions	Mean	
	Median	
perform_feature_selection	フラグ	

表 134. knnnode プロパティ (続き)

knnnode プロパティ	値	プロパティの説明
forced_entry_inputs	[field1 ... fieldN]	
stop_on_error_ratio	フラグ	
number_to_select	integer	
minimum_change	number	
validation_fold_assign_by_field	フラグ	
number_of_folds	integer	validation_fold_assign_by_field が False の場合にのみ有効です。
set_random_seed	フラグ	
random_seed	number	
folds_field	フィールド	validation_fold_assign_by_field が True の場合にのみ有効です。
all_probabilities	フラグ	
save_distances	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

kohonennode プロパティ



Kohonen ノードは、ニューラル・ネットワークの一種であり、データ・セットをクラスター化して異なるグループを形成する目的で使用できます。ネットワークの学習が完了すると、類似のレコードは出力マップで互い近くに表示され、違いの大きいレコードほど離れたところに表示されます。強度の高いユニットを識別するために生成されたモデル内で、各ユニットが獲得した観察の数値を調べることができます。これは、適切なクラスター数についてのヒントになる場合があります。

例

```
node = stream.create("kohonen", "My node")
# "Model" tab
node.setPropertyValue("use_model_name", False)
node.setPropertyValue("model_name", "Symbolic Cluster")
node.setPropertyValue("stop_on", "Time")
node.setPropertyValue("time", 1)
node.setPropertyValue("set_random_seed", True)
node.setPropertyValue("random_seed", 12345)
node.setPropertyValue("optimize", "Speed")
# "Expert" tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("width", 3)
node.setPropertyValue("length", 3)
node.setPropertyValue("decay_style", "Exponential")
```

```
node.setPropertyValue("phase1_neighborhood", 3)
node.setPropertyValue("phase1_eta", 0.5)
node.setPropertyValue("phase1_cycles", 10)
node.setPropertyValue("phase2_neighborhood", 1)
node.setPropertyValue("phase2_eta", 0.2)
node.setPropertyValue("phase2_cycles", 75)
```

表 135. kohonennode プロパティ

kohonennode プロパティ	値	プロパティの説明
inputs	[field1 ... fieldN]	Kohonen モデルは対象フィールドでなく、入力フィールドのリストを使用します。度数フィールドおよび重みフィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
continue	フラグ	
show_feedback	フラグ	
stop_on	Default Time	
time	number	
optimize	Speed Memory	モデル作成が速度とメモリのどちらにより最適化されるかを指定します。
cluster_label	フラグ	
mode	Simple Expert	
width	number	
length	number	
decay_style	Linear Exponential	
phase1_neighborhood	number	
phase1_eta	number	
phase1_cycles	number	
phase2_neighborhood	number	
phase2_eta	number	
phase2_cycles	number	

linearnode プロパティ



線型回帰モデルは、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。

例

```
node = stream.create("linear", "My node")
# Build Options tab - Objectives panel
node.setPropertyValue("objective", "Standard")
# Build Options tab - Model Selection panel
node.setPropertyValue("model_selection", "BestSubsets")
node.setPropertyValue("criteria_best_subsets", "ASE")
# Build Options tab - Ensembles panel
node.setPropertyValue("combining_rule_categorical", "HighestMeanProbability")
```

表 136. linearnode プロパティ

linearnode プロパティ	値	プロパティの説明
target	フィールド	1 つの対象フィールドを指定します。
inputs	[field1 ... fieldN]	モデルで使用される入力または入力または予測変数フィールド。
continue_training_existing_model	フラグ	
objective	Standard Bagging Boosting psm	psm は非常に大きいデータセットに使用され、Server の接続が必要です。
use_auto_data_preparation	フラグ	
confidence_level	number	
model_selection	ForwardStepwise BestSubsets None	
criteria_forward_stepwise	AICC Fstatistics AdjustedRSquare ASE	
probability_entry	number	

表 136. *linearnode* プロパティ (続き)

linearnode プロパティ	値	プロパティの説明
probability_removal	<i>number</i>	
use_max_effects	フラグ	
max_effects	<i>number</i>	
use_max_steps	フラグ	
max_steps	<i>number</i>	
criteria_best_subsets	AICC AdjustedRSquare ASE	
combining_rule_continuou s	Mean Median	
component_models_n	<i>number</i>	
use_random_seed	フラグ	
random_seed	<i>number</i>	
use_custom_model_name	フラグ	
custom_model_name	<i>string</i>	
use_custom_name	フラグ	
custom_name	<i>string</i>	
tooltip	<i>string</i>	
keywords	<i>string</i>	
annotation	<i>string</i>	

linearnode プロパティ



線型回帰モデルは、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。

表 137. *linearnode* プロパティ

linearnode プロパティ	値	プロパティの説明
target	フィールド	1 つの対象フィールドを指定します。
inputs	[<i>field1 ... fieldN</i>]	モデルで使用される入力または入力または予測変数フィールド。
weight_field	フィールド	モデルで使用される分析フィールド。
custom_fields	フラグ	デフォルト値は TRUE です。

表 137. linearasnode プロパティ (続き)

linearasnode プロパティ	値	プロパティの説明
intercept	フラグ	デフォルト値は TRUE です。
detect_2way_interaction	フラグ	2 要因の交互作用を考慮するかどうか。デフォルト値は TRUE です。
cin	<i>number</i>	モデル係数の推定値を計算するために使用する確信度の区間。0 より大きく 100 未満の値を指定してください。デフォルト値は 95 です。
factor_order	ascending descending	カテゴリ型予測フィールドの並び順。デフォルト値は ascending です。
var_select_method	ForwardStepwise BestSubsets none	使用するモデルの選択方法。デフォルト値は ForwardStepwise です。
criteria_for_forward_stepwise	AICC Fstatistics AdjustedRSquare ASE	モデルに効果を加えるべきか、またはモデルから効果を削除するべきかを決定するときに使用する統計。デフォルト値は AdjustedRSquare です。
pin	<i>number</i>	ここに指定された pin しきい値未満の最小 p 値を持つ効果がモデルに追加されます。デフォルト値は 0.05 です。
pout	<i>number</i>	ここに指定された pout しきい値より大きい p 値を持つモデル内のすべての効果が削除されます。デフォルト値は 0.10 です。
use_custom_max_effects	フラグ	最終モデルで最大数の効果を使用するかどうか。デフォルト値は FALSE です。
max_effects	<i>number</i>	最終モデルで使用する効果の最大数。デフォルト値は 1 です。
use_custom_max_steps	フラグ	最大数のステップを使用するかどうか。デフォルト値は FALSE です。
max_steps	<i>number</i>	ステップワイズ アルゴリズムが停止する最大ステップ数。デフォルト値は 1 です。

表 137. linearasnode プロパティ (続き)

linearasnode プロパティ	値	プロパティの説明
criteria_for_best_subsets	AICC	使用する基準のモード。デフォルト値は AdjustedRSquare です。
	AdjustedRSquare	
	ASE	

logregnode プロパティ



ロジスティック回帰は、入力フィールドの値に基づいてレコードを分類する統計手法です。線形回帰と似ていますが、数値範囲ではなくカテゴリ対象フィールドを使用します。

Multinomial Example

```
node = stream.create("logreg", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("target", "Drug")
node.setPropertyValue("inputs", ["BP", "Cholesterol", "Age"])
node.setPropertyValue("partition", "Test")
# "Model" tab
node.setPropertyValue("use_model_name", True)
node.setPropertyValue("model_name", "Log_reg Drug")
node.setPropertyValue("use_partitioned_data", True)
node.setPropertyValue("method", "Stepwise")
node.setPropertyValue("logistic_procedure", "Multinomial")
node.setPropertyValue("multinomial_base_category", "BP")
node.setPropertyValue("model_type", "FullFactorial")
node.setPropertyValue("custom_terms", [["BP", "Sex"], ["Age"], ["Na", "K"]])
node.setPropertyValue("include_constant", False)
# "Expert" tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("scale", "Pearson")
node.setPropertyValue("scale_value", 3.0)
node.setPropertyValue("all_probabilities", True)
node.setPropertyValue("tolerance", "1.0E-7")
# "Convergence..." section
node.setPropertyValue("max_iterations", 50)
node.setPropertyValue("max_steps", 3)
node.setPropertyValue("l_converge", "1.0E-3")
node.setPropertyValue("p_converge", "1.0E-7")
node.setPropertyValue("delta", 0.03)
# "Output..." section
node.setPropertyValue("summary", True)
node.setPropertyValue("likelihood_ratio", True)
node.setPropertyValue("asymptotic_correlation", True)
node.setPropertyValue("goodness_fit", True)
node.setPropertyValue("iteration_history", True)
node.setPropertyValue("history_steps", 3)
node.setPropertyValue("parameters", True)
node.setPropertyValue("confidence_interval", 90)
node.setPropertyValue("asymptotic_covariance", True)
node.setPropertyValue("classification_table", True)
# "Stepping" options
node.setPropertyValue("min_terms", 7)
node.setPropertyValue("use_max_terms", True)
node.setPropertyValue("max_terms", 10)
```

```
node.setPropertyValue("probability_entry", 3)
node.setPropertyValue("probability_removal", 5)
node.setPropertyValue("requirements", "Containment")
```

Binomial Example

```
node = stream.create("logreg", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("target", "Cholesterol")
node.setPropertyValue("inputs", ["BP", "Drug", "Age"])
node.setPropertyValue("partition", "Test")
# "Model" tab
node.setPropertyValue("use_model_name", False)
node.setPropertyValue("model_name", "Log_reg Cholesterol")
node.setPropertyValue("multinomial_base_category", "BP")
node.setPropertyValue("use_partitioned_data", True)
node.setPropertyValue("binomial_method", "Forwards")
node.setPropertyValue("logistic_procedure", "Binomial")
node.setPropertyValue("binomial_categorical_input", "Sex")
node.setKeyedPropertyValue("binomial_input_contrast", "Sex", "Simple")
node.setKeyedPropertyValue("binomial_input_category", "Sex", "Last")
node.setPropertyValue("include_constant", False)
# "Expert" tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("scale", "Pearson")
node.setPropertyValue("scale_value", 3.0)
node.setPropertyValue("all_probabilities", True)
node.setPropertyValue("tolerance", "1.0E-7")
# "Convergence..." section
node.setPropertyValue("max_iterations", 50)
node.setPropertyValue("l_converge", "1.0E-3")
node.setPropertyValue("p_converge", "1.0E-7")
# "Output..." section
node.setPropertyValue("binomial_output_display", "at_each_step")
node.setPropertyValue("binomial_goodness_of_fit", True)
node.setPropertyValue("binomial_iteration_history", True)
node.setPropertyValue("binomial_parameters", True)
node.setPropertyValue("binomial_ci_enable", True)
node.setPropertyValue("binomial_ci", 85)
# "Stepping" options
node.setPropertyValue("binomial_removal_criterion", "LR")
node.setPropertyValue("binomial_probability_removal", 0.2)
```

表 138. logregnode プロパティ

logregnode プロパティ	値	プロパティの説明
target	フィールド	ロジスティック回帰モデルは1つの対象フィールドおよび1つ以上の入力フィールドを使用します。度数フィールドおよび重みフィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
logistic_procedure	Binomial Multinomial	
include_constant	フラグ	

表 138. logregnode プロパティ (続き)

logregnode プロパティ	値	プロパティの説明
mode	Simple Expert	
method	Enter Stepwise Forwards Backwards BackwardsStepwise	
binomial_method	Enter Forwards Backwards	
model_type	MainEffects FullFactorial Custom	モデル・タイプとして FullFactorial が指定されている場合、ステップ手法が指定されたとしても、実行されません。その代わりに、強制投入法 (Enter) が使用されます。 モデル・タイプに Custom が設定されてもユーザー設定フィールド (custom fields) が指定されていない場合は、主効果モデルが構築されます。
custom_terms	[[BP Sex][BP][Age]]	
multinomial_base_category	string	参照カテゴリーの決定方法を指定します。
binomial_categorical_input	string	
binomial_input_contrast	Indicator Simple Difference Helmert Repeated Polynomial Deviation	コントラストを決定する方法を指定するカテゴリー入力用のキー・プロパティ。

表 138. logregnode プロパティ (続き)

logregnode プロパティ	値	プロパティの説明
binomial_input_category	First Last	参照カテゴリーを決定する方法を指定するカテゴリー入力用のキー・プロパティ。
scale	None UserDefined Pearson Deviance	
scale_value	<i>number</i>	
all_probabilities	フラグ	
tolerance	1.0E-5 1.0E-6 1.0E-7 1.0E-8 1.0E-9 1.0E-10	
min_terms	<i>number</i>	
use_max_terms	フラグ	
max_terms	<i>number</i>	
entry_criterion	Score LR	
removal_criterion	LR Wald	
probability_entry	<i>number</i>	
probability_removal	<i>number</i>	
binomial_probability_entry	<i>number</i>	
binomial_probability_removal	<i>number</i>	

表 138. logregnode プロパティ (続き)

logregnode プロパティ	値	プロパティの説明
requirements	HierarchyDiscrete HierarchyAll Containment None	
max_iterations	<i>number</i>	
max_steps	<i>number</i>	
p_converge	1.0E-4 1.0E-5 1.0E-6 1.0E-7 1.0E-8 0	
l_converge	1.0E-1 1.0E-2 1.0E-3 1.0E-4 1.0E-5 0	
delta	<i>number</i>	
iteration_history	フラグ	
history_steps	<i>number</i>	
summary	フラグ	
likelihood_ratio	フラグ	
asymptotic_correlation	フラグ	
goodness_fit	フラグ	
parameters	フラグ	
confidence_interval	<i>number</i>	
asymptotic_covariance	フラグ	
classification_table	フラグ	

表 138. logregnode プロパティ (続き)

logregnode プロパティ	値	プロパティの説明
stepwise_summary	フラグ	
info_criteria	フラグ	
monotonicity_measures	フラグ	
binomial_output_display	at_each_step at_last_step	
binomial_goodness_of_fit	フラグ	
binomial_parameters	フラグ	
binomial_iteration_history	フラグ	
binomial_classification_plots	フラグ	
binomial_ci_enable	フラグ	
binomial_ci	number	
binomial_residual	outliers all	
binomial_residual_enable	フラグ	
binomial_outlier_threshold	number	
binomial_classification_cutoff	number	
binomial_removal_criterion	LR Wald Conditional	
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	

lsvmnode プロパティ



線型サポート・ベクター・マシン (LSVM) ノードを使用すると、オーバーフィットすることなく、データを2つのグループのいずれかに分類することができます。LSVMは線型であり、極めて多数のレコードを含むデータセットなど、広範なデータセットを処理することができます。

表 139. *lsvmnode* プロパティ

lsvmnode プロパティ	値	プロパティの説明
intercept	フラグ	モデルに切片を含めます。デフォルト値は True です。
target_order	Ascending Descending	カテゴリ型対象のソート順を指定します。連続型対象では無視されます。デフォルトは Ascending です。
precision	number	対象フィールドの尺度が Continuous の場合にのみ使用されます。回帰の損失の感度に関連するパラメーターを指定します。最小値は 0 で最大値はありません。デフォルト値は 0.1 です。
exclude_missing_values	フラグ	True にすると、欠損値が 1 つでもある場合はレコードが除外されます。デフォルト値は False です。
penalty_function	L1 L2	使用するペナルティ関数のタイプを指定します。デフォルト値は L2 です。
lambda	number	ペナルティ (正規化) パラメーター。
calculate_variable_importance	フラグ	重要度の適切な測定を生成するモデルの場合、このオプションにより、モデルの推定における各予測値の相対重要度を示すグラフが表示されます。一部のモデルでは、特に大規模データセットを処理する場合、変数の重要度の計算には長い時間がかかるため、一部のモデルではデフォルトでオフになっています。変数の重要度は、ディシジョンリストモデルでは使用できません。

neuralnetnode プロパティ

重要: 機能が拡張された新しいバージョンのニューラル・ネットワーク・モデル作成ノードがこのリリースで使用できます。新しいバージョンについては次の項で説明します (*neuralnetwork*)。旧バージョンでモデルを作成およびスコアリングできますが、新しいバージョンを使用するようスクリプトを更新することをお勧めします。以下は旧バージョンの詳細です。

例

```
node = stream.create("neuralnet", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("targets", ["Drug"])
node.setPropertyValue("inputs", ["Age", "Na", "K", "Cholesterol", "BP"])
# "Model" tab
node.setPropertyValue("use_partitioned_data", True)
node.setPropertyValue("method", "Dynamic")
node.setPropertyValue("train_pct", 30)
```

```

node.setPropertyValue("set_random_seed", True)
node.setPropertyValue("random_seed", 12345)
node.setPropertyValue("stop_on", "Time")
node.setPropertyValue("accuracy", 95)
node.setPropertyValue("cycles", 200)
node.setPropertyValue("time", 3)
node.setPropertyValue("optimize", "Speed")
# "Multiple Method Expert Options" section
node.setPropertyValue("m_topologies", "5 30 5; 2 20 3, 1 10 1")
node.setPropertyValue("m_non_pyramids", False)
node.setPropertyValue("m_persistence", 100)

```

表 140. *neuralnetnode* プロパティ

neuralnetnode プロパティ	値	プロパティの説明
targets	[<i>field1 ... fieldN</i>]	ニューラル・ネット・ノードには、1つ以上の対象フィールドと1つ以上の入力フィールドが必要です。度数および重みフィールドは無視されます。詳しくは、トピック 217 ページの『 一般的なモデル作成ノードのプロパティ 』を参照してください。
method	Quick Dynamic Multiple Prune ExhaustivePrune RBFN	
prevent_overtrain	フラグ	
train_pct	<i>number</i>	
set_random_seed	フラグ	
random_seed	<i>number</i>	
mode	Simple Expert	
stop_on	Default Accuracy Cycles Time	停止モード。
accuracy	<i>number</i>	停止精度。
cycles	<i>number</i>	学習サイクル。

表 140. *neuralnetnode* プロパティ (続き)

neuralnetnode プロパティ	値	プロパティの説明
time	<i>number</i>	学習時間 (分)。
continue	フラグ	
show_feedback	フラグ	
binary_encode	フラグ	
use_last_model	フラグ	
gen_logfile	フラグ	
logfile_name	<i>string</i>	
alpha	<i>number</i>	
initial_eta	<i>number</i>	
high_eta	<i>number</i>	
low_eta	<i>number</i>	
eta_decay_cycles	<i>number</i>	
hid_layers	One Two Three	
hl_units_one	<i>number</i>	
hl_units_two	<i>number</i>	
hl_units_three	<i>number</i>	
persistence	<i>number</i>	
m_topologies	<i>string</i>	
m_non_pyramids	フラグ	
m_persistence	<i>number</i>	
p_hid_layers	One Two Three	
p_hl_units_one	<i>number</i>	
p_hl_units_two	<i>number</i>	
p_hl_units_three	<i>number</i>	
p_persistence	<i>number</i>	
p_hid_rate	<i>number</i>	
p_hid_pers	<i>number</i>	
p_inp_rate	<i>number</i>	

表 140. *neuralnetnode* プロパティ (続き)

neuralnetnode プロパティ	値	プロパティの説明
p_inp_pers	number	
p_overall_pers	number	
r_persistence	number	
r_num_clusters	number	
r_eta_auto	フラグ	
r_alpha	number	
r_eta	number	
optimize	Speed Memory	モデル作成が速度とメモリーのどちらにより最適化されるかを指定します。
calculate_variable_importance	フラグ	注: 前回のリリースで使った sensitivity_analysis プロパティは、このプロパティにより廃止されます。古いプロパティはまだサポートされますが、 calculate_variable_importance をお勧めします。
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

neuralnetworknode プロパティ



ニューラル・ネット・ノードは、人間の脳が情報を処理する方法を単純化したモデルを使用します。ニューラル・ネットワーク・ノードは、関係する多数の単純な処理単位をシミュレートします。処理単位は、ニューロンを抽象化したものと表現できます。ニューラル・ネットワークは強力な一般関数推定法であり、学習させたり、適用するには、最低限の統計学および数学の知識しか必要ありません。

例

```
node = stream.create("neuralnetwork", "My node")
# Build Options tab - Objectives panel
node.setPropertyValue("objective", "Standard")
# Build Options tab - Ensembles panel
node.setPropertyValue("combining_rule_categorical", "HighestMeanProbability")
```

表 141. *neuralnetworknode* プロパティ

neuralnetworknode プロパティ	値	プロパティの説明
targets	[field1 ... fieldN]	対象フィールドを指定します。

表 141. *neuralnetworknode* プロパティ (続き)

neuralnetworknode プロパティ	値	プロパティの説明
inputs	[<i>field1</i> ... <i>fieldN</i>]	モデルで使用される入力または入力または予測変数フィールド。
splits	[<i>field1</i> ... <i>fieldN</i>]	分割モデル作成に使用する、フィールドを選択します。
use_partition	フラグ	区分フィールドが定義される場合、このオプションは学習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。
continue	フラグ	既存モデルの学習を継続:
objective	Standard Bagging Boosting psm	psm は非常に大きいデータセットに使用され、Server の接続が必要です。
method	MultilayerPerceptron RadialBasisFunction	
use_custom_layers	フラグ	
first_layer_units	<i>number</i>	
second_layer_units	<i>number</i>	
use_max_time	フラグ	
max_time	<i>number</i>	
use_max_cycles	フラグ	
max_cycles	<i>number</i>	
use_min_accuracy	フラグ	
min_accuracy	<i>number</i>	
combining_rule_categorical	Voting HighestProbability HighestMeanProbability	
combining_rule_continuou s	Mean Median	
component_models_n	<i>number</i>	
overfit_prevention_pct	<i>number</i>	

表 141. *neuralnetworknode* プロパティ (続き)

neuralnetworknode プロパティ	値	プロパティの説明
use_random_seed	フラグ	
random_seed	<i>number</i>	
missing_values	listwiseDeletion missingValueImputation	
use_model_name	<i>Boolean</i>	
model_name	<i>string</i>	
confidence	onProbability onIncrease	
score_category_probabilities	フラグ	
max_categories	<i>number</i>	
score_propensity	フラグ	
use_custom_name	フラグ	
custom_name	<i>string</i>	
tooltip	<i>string</i>	
keywords	<i>string</i>	
annotation	<i>string</i>	

questnode プロパティ



QUEST ノードには、ディシジョン・ツリーの構築用に 2 分岐の方法が用意されています。これは、大規模な C&R Tree 分析が必要とする処理時間を短縮すると同時に、より多くの分割を可能にする入力値が優先される分類ツリー内の傾向を低減するように設計されています。入力フィールドは、数値範囲 (連続型) にできますが、目標変数はカテゴリーでなければなりません。すべての分割は 2 進数です。

例

```
node = stream.create("quest", "My node")
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("target", "Drug")
node.setPropertyValue("inputs", ["Age", "Na", "K", "Cholesterol", "BP"])
node.setPropertyValue("model_output_type", "InteractiveBuilder")
node.setPropertyValue("use_tree_directives", True)
node.setPropertyValue("max_surrogates", 5)
node.setPropertyValue("split_alpha", 0.03)
node.setPropertyValue("use_percentage", False)
node.setPropertyValue("min_parent_records_abs", 40)
node.setPropertyValue("min_child_records_abs", 30)
node.setPropertyValue("prune_tree", True)
node.setPropertyValue("use_std_err", True)
node.setPropertyValue("std_err_multiplier", 3)
```

表 142. questnode プロパティ

questnode プロパティ	値	プロパティの説明
target	フィールド	QUEST モデルは単一の対象フィールドおよび 1 つ以上の入力フィールドを使用します。度数フィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
continue_training_existing_model	フラグ	
objective	Standard Boosting Bagging psm	psm は非常に大きいデータセットに使用され、Server の接続が必要です。
model_output_type	Single InteractiveBuilder	
use_tree_directives	フラグ	
tree_directives	string	
use_max_depth	Default Custom	
max_depth	integer	最大ツリー深さ (0 から 1000)。use_max_depth = Custom の場合にのみ使用されます。
prune_tree	フラグ	オーバーフィットしないようにツリーを剪定します。
use_std_err	フラグ	リスクにおける最大差 (標準誤差) を使用します。
std_err_multiplier	number	最大差。
max_surrogates	number	最大代理変数。
use_percentage	フラグ	
min_parent_records_pc	number	
min_child_records_pc	number	
min_parent_records_abs	number	
min_child_records_abs	number	
use_costs	フラグ	
costs	構造化	構造化プロパティ。

表 142. *questnode* プロパティ (続き)

questnode プロパティ	値	プロパティの説明
priors	Data Equal Custom	
custom_priors	構造化	構造化プロパティ。
adjust_priors	フラグ	
trails	<i>number</i>	ブーストまたはバグのコンポーネント・モデル数。
set_ensemble_method	Voting HighestProbability HighestMeanProbability	カテゴリー型対象のデフォルト結合ルール。
range_ensemble_method	Mean Median	連続型対象のデフォルト結合ルール。
large_boost	フラグ	特に大きなデータセットのブースティングを適用します。
split_alpha	<i>number</i>	分割の有意水準。
train_pct	<i>number</i>	オーバーフィット防止セット。
set_random_seed	フラグ	結果を再現オプション。
seed	<i>number</i>	
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

randomtrees プロパティ



Random Trees ノードは既存の C & R Tree ノードと似ていますが、Random Trees ノードはビッグデータを処理して単一のツリーを作成するように設計されており、結果のモデルが スポス モデラー バージョン 17 で追加された出力ビューアーに表示されます。Random Trees ノードが生成するディジションツリーを使用して、将来の観測値を予測または分類できます。この方法では、再帰的なデータ分岐を使用して、各ステップで不純性を最小限に抑えることで、学習レコードがセグメントに分割されます。ツリー内のノードは、ノード内のケースの 100% が対象フィールドの特定のカテゴリに分類される場合に、純粋 と見なされます。対象フィールドおよび入力フィールドは、数値範囲またはカテゴリ (名義型、順序型、フラグ) が使用できます。すべての分岐は 2 分割です (2 つのサブグループのみ)。

表 143. randomtrees プロパティ

randomtrees プロパティ	値	プロパティの説明
target	フィールド	Random Trees ノードでは、モデルには単一の対象フィールドおよび 1 つ以上の入力フィールドが必要になります。度数フィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
number_of_models	integer	アンサンブル・モデル構築の一環として構築されるモデルの数を決定します。
use_number_of_predictors	フラグ	number_of_predictors を使用するかどうかを決定します。
number_of_predictors	integer	分割モデルの構築時に使用する予測値の個数を指定します。
use_stop_rule_for_accuracy	フラグ	精度を向上できない場合にモデル構築を中止するかどうかを決定します。
sample_size	number	極めて大規模なデータ・セットを処理する際にパフォーマンスを向上させるには、この値を小さくします。
handle_imbalanced_data	フラグ	モデルの対象が特定のフラグの結果であり、望ましくない結果に対する望まれる結果の比率が非常に小さい場合は、データは不均衡になり、モデルによって実行されるブートストラップ・サンプリングがモデルの精度に影響を与える可能性があります。不均衡なデータの処理を有効にすると、モデルが収集する望ましい結果の比率が高まり、より強固なモデルが生成されます。
use_weighted_sampling	フラグ	False の場合、各ノードの変数は、同じ確率で無作為に選択されます。True の場合、変数には重みが付けられ、それに応じて選択されます。

表 143. *randomtrees* プロパティ (続き)

randomtrees プロパティ	値	プロパティの説明
max_node_number	<i>integer</i>	個々のツリーで許容されるノードの最大数。次の分割でこの数を超えることが予想される場合、ツリーの成長は停止します。
max_depth	<i>integer</i>	ツリーの最大の深さ。これに達すると成長は停止します。
min_child_node_size	<i>integer</i>	親ノードの分割後に子ノードで許容されるレコードの最小数を決定します。子ノードに含まれることになるレコードの数がここで指定した数よりも少ない場合、親ノードは分割されません。
use_costs	フラグ	
costs	構造化	構造化プロパティ。形式は、実際の値、予測された値、およびコスト (予測が正しくない場合) の 3 つの値のリストです。例： <code>tree.setPropertyValue("costs", [["drugA", "drugB", 3.0], ["drugX", "drugY", 4.0]])</code>
default_cost_increase	none linear square custom	注: 順序型対象に対してのみ有効です。 コスト行列にデフォルト値を設定します。
max_pct_missing	<i>integer</i>	いずれかの入力の欠損値の割合がここで指定した値より大きい場合、その入力は除外されます。最小値は 0、最大値は 100 です。
exclude_single_cat_pct	<i>integer</i>	いずれかのカテゴリー値がここで指定したレコードの割合より高い場合、そのフィールド全体がモデル構築から除外されます。最小値は 1、最大値は 99 です。
max_category_number	<i>integer</i>	フィールド内のカテゴリー数がこの値を超える場合、そのフィールドはモデル構築から除外されます。最大値は 2 です。
min_field_variation	<i>number</i>	連続型フィールドの変動係数がこの値より小さい場合、そのフィールドはモデル構築から除外されます。

表 143. *randomtrees* プロパティ (続き)

randomtrees プロパティ	値	プロパティの説明
num_bins	<i>integer</i>	データが連続型入力で構成される場合にのみ使用されます。入力に対して使用する等しいフリクエンシ ビンの数を設定します。オプションは 2、4、5、10、20、25、50、または 100 です。
topN	<i>integer</i>	報告するルール数を指定します。デフォルト値は 50 で、最小値は 1、最大値は 1000 です。

regressionnode プロパティ



線形回帰は、データを要約し、予測された出力値と実際の出力値の間の矛盾を最小化する直線または面を当てはめることによって予測を行うための一般的な統計手法です。

注: 今後のリリースでは、線形回帰ノードは線型ノードに置き換えられる予定になっています。今後、線形回帰には線型モデルを使用することをお勧めします。

例

```
node = stream.create("regression", "My node")
# "Fields" tab
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("target", "Age")
node.setPropertyValue("inputs", ["Na", "K"])
node.setPropertyValue("partition", "Test")
node.setPropertyValue("use_weight", True)
node.setPropertyValue("weight_field", "Drug")
# "Model" tab
node.setPropertyValue("use_model_name", True)
node.setPropertyValue("model_name", "Regression Age")
node.setPropertyValue("use_partitioned_data", True)
node.setPropertyValue("method", "Stepwise")
node.setPropertyValue("include_constant", False)
# "Expert" tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("complete_records", False)
node.setPropertyValue("tolerance", "1.0E-3")
# "Stepping..." section
node.setPropertyValue("stepping_method", "Probability")
node.setPropertyValue("probability_entry", 0.77)
node.setPropertyValue("probability_removal", 0.88)
node.setPropertyValue("F_value_entry", 7.0)
node.setPropertyValue("F_value_removal", 8.0)
# "Output..." section
node.setPropertyValue("model_fit", True)
node.setPropertyValue("r_squared_change", True)
node.setPropertyValue("selection_criteria", True)
node.setPropertyValue("descriptives", True)
node.setPropertyValue("p_correlations", True)
node.setPropertyValue("collinearity_diagnostics", True)
node.setPropertyValue("confidence_interval", True)
node.setPropertyValue("covariance_matrix", True)
node.setPropertyValue("durbin_watson", True)
```

表 144. regressionnode プロパティ

regressionnode プロパティ	値	プロパティの説明
target	フィールド	回帰モデルは単一の対象フィールドおよび 1 つ以上の入力フィールドを使用します。重みフィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
method	Enter Stepwise Backwards Forwards	
include_constant	フラグ	
use_weight	フラグ	
weight_field	フィールド	
mode	Simple Expert	
complete_records	フラグ	
tolerance	1.0E-1 1.0E-2 1.0E-3 1.0E-4 1.0E-5 1.0E-6 1.0E-7 1.0E-8 1.0E-9 1.0E-10 1.0E-11 1.0E-12	引数には二重引用符を使用します。

表 144. regressionnode プロパティ (続き)

regressionnode プロパティ	値	プロパティの説明
stepping_method	useP	useP : F 値確率を使用
	useF	useF: F 値を使用
probability_entry	number	
probability_removal	number	
F_value_entry	number	
F_value_removal	number	
selection_criteria	フラグ	
confidence_interval	フラグ	
covariance_matrix	フラグ	
collinearity_diagnostics	フラグ	
regression_coefficients	フラグ	
exclude_fields	フラグ	
durbin_watson	フラグ	
model_fit	フラグ	
r_squared_change	フラグ	
p_correlations	フラグ	
descriptives	フラグ	
calculate_variable_importance	フラグ	

sequencenode プロパティ



シーケンス・ノードで、シーケンシャルな、または時間経過が伴うデータ内のアソシエーション・ルールを検出します。シーケンスは、予測可能な順序で発生する傾向があるアイテム・セットのリストです。例えば、ラゾルやアフターシェーブローションを購入する顧客は、次回の店ではシェービングクリームを購入することがあります。シーケンス・ノードは CARMA アソシエーション・ルール・アルゴリズムに基づいており、効率的な 2 段階通過法を使用してシーケンスを検出します。

例

```
node = stream.create("sequence", "My node")
# "Fields" tab
node.setPropertyValue("id_field", "Age")
node.setPropertyValue("contiguous", True)
node.setPropertyValue("use_time_field", True)
node.setPropertyValue("time_field", "Date1")
node.setPropertyValue("content_fields", ["Drug", "BP"])
node.setPropertyValue("partition", "Test")
# "Model" tab
node.setPropertyValue("use_model_name", True)
node.setPropertyValue("model_name", "Sequence_test")
node.setPropertyValue("use_partitioned_data", False)
node.setPropertyValue("min_supp", 15.0)
```

```

node.setPropertyValue("min_conf", 14.0)
node.setPropertyValue("max_size", 7)
node.setPropertyValue("max_predictions", 5)
# "Expert" tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("use_max_duration", True)
node.setPropertyValue("max_duration", 3.0)
node.setPropertyValue("use_pruning", True)
node.setPropertyValue("pruning_value", 4.0)
node.setPropertyValue("set_mem_sequences", True)
node.setPropertyValue("mem_sequences", 5.0)
node.setPropertyValue("use_gaps", True)
node.setPropertyValue("min_item_gap", 20.0)
node.setPropertyValue("max_item_gap", 30.0)

```

表 145. sequencenode プロパティ

sequencenode プロパティ	値	プロパティの説明
id_field	フィールド	シーケンス・モデルを作成するには、ID フィールドを指定する必要があります。さらにオプションで時間フィールドと 1 つ以上の内容フィールドを指定します。重みフィールドおよび度数フィールドは使用しません。詳しくは、トピック 217 ページの『 一般的なモデル作成ノードのプロパティ 』を参照してください。
time_field	フィールド	
use_time_field	フラグ	
content_fields	[field1 ... fieldN]	
contiguous	フラグ	
min_supp	number	
min_conf	number	
max_size	number	
max_predictions	number	
mode	Simple Expert	
use_max_duration	フラグ	
max_duration	number	
use_gaps	フラグ	
min_item_gap	number	
max_item_gap	number	
use_pruning	フラグ	
pruning_value	number	
set_mem_sequences	フラグ	
mem_sequences	integer	

slrmnode プロパティ



SLRM (自己学習応答モデル) ノードを使用するとモデルを構築でき、単一または少数の新しいケースを使用して全データを使用するモデルの保持をすることなく、モデルの再見積もりを行うことができます。

例

```
node = stream.create("slrm", "My node")
node.setPropertyValue("target", "Offer")
node.setPropertyValue("target_response", "Response")
node.setPropertyValue("inputs", ["Cust_ID", "Age", "Ave_Bal"])
```

表 146. slrmnode プロパティ

slrmnode プロパティ	値	プロパティの説明
target	フィールド	対象フィールドは名義型またはフラグ型である必要があります。度数フィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
target_response	フィールド	フラグ型である必要があります。
continue_training_existing_model	フラグ	
target_field_values	フラグ	すべて使用:ソースのすべての値を使用します。 指定:必要な値を選択します。
target_field_values_specify	[field1 ... fieldN]	
include_model_assessment	フラグ	
model_assessment_random_seed	number	実数である必要があります。
model_assessment_sample_size	number	実数である必要があります。
model_assessment_iterations	number	反復数。
display_model_evaluation	フラグ	
max_predictions	number	
randomization	number	
scoring_random_seed	number	
sort	Ascending Descending	高いスコアまたは低いスコアのどちらを持つオファーが最初に表示されるかを指定します。
model_reliability	フラグ	

表 146. <i>slrmnode</i> プロパティ (続き)		
slrmnode プロパティ	値	プロパティの説明
calculate_variable_importance	フラグ	

statisticsmodelnode プロパティ



Statistics モデル・ノードを使用すると、PMML を作成する IBM SPSS 統計 手続きを実行してデータを分析および使用することができます。このノードには、ライセンス交付を受けた IBM SPSS 統計のコピーが必要です。

このノードのプロパティについては、[424 ページの『statisticsmodelnode プロパティ』](#)に記載されています。

stptime プロパティ



時空間予測 (STP) ノードは、ロケーション・データ、予測用の入力フィールド (予測値)、時間フィールド、および対象フィールドを使用します。各ロケーションには、それぞれの測定時の各予測値を表すデータの行が多数あります。データを分析すると、そのデータを使用して、分析で使用する形状データ内の任意のロケーションの対象値を予測できます。

表 147. <i>stptime</i> プロパティ		
stptime プロパティ	データ・タイプ	プロパティの説明
「フィールド」タブ		
target	フィールド	これは対象フィールドです。
location	フィールド	モデルの場所フィールド。地理空間フィールドのみ許可されます。
location_label	フィールド	location で選択された場所にラベルを付けるために出力内で使用されるカテゴリ型フィールド。
time_field	フィールド	モデルの時間フィールド。連続型の尺度を持つフィールドのみ許可されます。ストレージタイプは、時間、日付、タイムスタンプ、整数のいずれかでなければなりません。
inputs	[field1 ... fieldN]	入力フィールドのリスト。
「時間区分」タブ		

表 147. stpnode プロパティ (続き)

stpnode プロパティ	データ・タイプ	プロパティの説明
interval_type_timestamp	Years Quarters Months Weeks Days Hours Minutes Seconds	
interval_type_date	Years Quarters Months Weeks Days	
interval_type_time	Hours Minutes Seconds	STP が計算で使用する時間インデックスの作成時に処理対象となる週あたりの日数を制限します。
interval_type_integer	Periods (時間インデックス フィールドの場合のみ、整数のストレージ)	データ セットを変換する間隔。選択できる項目は、モデルの time_field として選択されたフィールドのストレージタイプによって異なります。
period_start	integer	

表 147. *stpnode* プロパティ (続き)

stpnode プロパティ	データ・タイプ	プロパティの説明
start_month	January February March April May June July August September October November December	モデルがインデックス作成を開始する月です。例えば、March に設定した場合、データセットの最初のレコードが January であるとしたら、モデルは最初の 2 つのレコードをスキップして 3 月からインデックス作成を開始します。
week_begins_on	Sunday Monday Tuesday Wednesday Thursday Friday Saturday	STP がデータから作成した時間インデックスの開始点。
days_per_week	<i>integer</i>	最小値は 1、最大値は 7、増分値は 1 です。
hours_per_day	<i>integer</i>	1 日のうちで、そのモデルが占める時間数。例えば、10 に設定した場合、モデルは day_begins_at の時刻に開始され、10 時間にわたってインデックス作成を続け、day_begins_at 値に一致する次の値までスキップします。

表 147. stpnode プロパティ (続き)		
stpnode プロパティ	データ・タイプ	プロパティの説明
day_begins_at	00:00 01:00 02:00 03:00 ... 23:00	モデルがインデックス作成を開始する時間の値を設定します。
interval_increment	1 2 3 4 5 6 10 12 15 20 30	この増分の設定は分または秒に対応します。これは、モデルがデータのインデックス作成を開始する位置を決定します。つまり、増分が 30 で間隔のタイプが seconds の場合、モデルはデータのインデックス作成を 30 秒ごとに行います。

表 147. *stpnode* プロパティ (続き)

stpnode プロパティ	データ・タイプ	プロパティの説明
data_matches_interval	Boolean	これを N に設定すると、モデルの構築前に、データが通常の interval_type に変換されます。 現在のデータがすでに正しい形式になっていて、interval_type とそれに関連するすべての設定がデータに一致している場合は、データの変換や集計が実行されないように、このプロパティを Y に設定してください。 このプロパティを Y に設定すると、すべての集計コントロールが無効になります。
agg_range_default	Sum Mean Min Max Median 1stQuartile 3rdQuartile	これは、連続型フィールドに使用されるデフォルトの集計方法を指定します。ユーザー指定の集計に明確に含まれていない連続型フィールドは、ここに指定した方法で集計されます。

表 147. stpnode プロパティ (続き)

stpnode プロパティ	データ・タイプ	プロパティの説明
custom_agg	<pre>[[field, aggregation method],[...]]</pre> デモ: <pre>[['x5' 'FirstQuartile']['x4' 'Sum']]</pre>	構造化プロパティ: スクリプト パラメーター: custom_agg 以下に例を示します。 <pre>set :stpnode.custom_agg = [[field1 function] [field2 function]]</pre> ここで、function は、当該フィールドで使用される集計関数です。
「基本」タブ		
include_intercept	フラグ	
max_autoregressive_lag	integer	1 の増分で最小 1、最大 5 です。 これは、予測に必要な以前のレコードの数です。したがって、例えば 5 に設定した場合は、以前の 5 件のレコードを使用して新しい予測が作成されます。ここに指定した、ビルド データからのレコード件数は、モデルに組み込まれます。したがって、ユーザーはモデルのスコアリング時にデータを再度提供する必要がありません。
estimation_method	Parametric Nonparametric	空間共分散行列のモデリング方法。
parametric_model	Gaussian Exponential PoweredExponential	Parametric 空間共分散モデルの順序パラメータ。
exponential_power	number	PoweredExponential モデルのべき乗レベル。最小値は 1、最大値は 2 です。
「詳細」タブ		

表 147. stpnode プロパティ (続き)

stpnode プロパティ	データ・タイプ	プロパティの説明
max_missing_values	integer	モデル内で許可される、欠損値を持つレコードの最大パーセント値。
significance	number	モデル構築における仮説検証の有意水準。STP モデル推定のすべての検定 (2 つの適合度検定、効果 F 検定、係数 T 検定を含む) に使用する有意水準値を指定します。
「出力」タブ		
model_specifications	フラグ	
temporal_summary	フラグ	
location_summary	フラグ	場所の要約表がモデル出力に含まれるかどうかを指定します。
model_quality	フラグ	
test_mean_structure	フラグ	
mean_structure_coefficients	フラグ	
autoregressive_coefficients	フラグ	
test_decay_space	フラグ	
parametric_spatial_covariance	フラグ	
correlations_heat_map	フラグ	
correlations_map	フラグ	
location_clusters	フラグ	
similarity_threshold	number	類似度のしきい値。この値を超えると、出力クラスターの類似度が十分に高いと判断され、1 つのクラスターに結合されます。
max_number_clusters	integer	モデル出力に含めることができるクラスターの上限值。
「モデル オプション」タブ		
use_model_name	フラグ	
model_name	string	
uncertainty_factor	number	最小値は 0、最大値は 100 です。将来の予測に適用される不確実性 (誤差) の増加を指定します。これは、予測の上限と下限です。

svmnode プロパティ



サポート・ベクター・マシン (SVM) ノードを使用すると、オーバーフィットすることなく、データを2つのグループのいずれかに分類することができます。SVMは、非常に多数の入力フィールドを含むデータセットなど、広範なデータセットを処理することができます。

例

```
node = stream.create("svm", "My node")
# Expert tab
node.setPropertyValue("mode", "Expert")
node.setPropertyValue("all_probabilities", True)
node.setPropertyValue("kernel", "Polynomial")
node.setPropertyValue("gamma", 1.5)
```

表 148. svmnode プロパティ

svmnode プロパティ	値	プロパティの説明
all_probabilities	フラグ	
stopping_criteria	1.0E-1 1.0E-2 1.0E-3 (default) 1.0E-4 1.0E-5 1.0E-6	最適化アルゴリズムをいつ停止するかを決定します。
regularization	number	C パラメーターとしても知られています。
precision	number	対象フィールドの尺度が Continuous の場合にのみ使用されます。
kernel	RBF (デフォルト) Polynomial Sigmoid Linear	変換に使用されるカーネル関数のタイプ。
rbf_gamma	number	kernel が RBF の場合にのみ使用されます。
gamma	number	kernel が Polynomial または Sigmoid の場合にのみ使用されます。
bias	number	

表 148. svmnode プロパティ (続き)

svmnode プロパティ	値	プロパティの説明
degree	number	kernel が Polynomial の場合にのみ使用されます。
calculate_variable_importance	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	
adjusted_propensity_partition	Test Validation	

tcmnode プロパティ



時間的因果モデリングでは、時系列データ内の重要な因果関係の検出が試行されます。時間的因果モデリングでは、一連の対象系列を指定し、それらの対象系列に対する一連の入力候補を指定します。その後、プロシーチャーは、各対象系列について自己回帰の時系列モデルを構築し、対象系列との重要な因果関係を持つ入力だけを取り込みます。

表 149. tcmnode プロパティ

tcmnode プロパティ	値	プロパティの説明
custom_fields	Boolean	
dimensionlist	[dimension1 ... dimensionN]	
data_struct	Multiple	
	Single	
metric_fields	フィールド	
both_target_and_input	[f1 ... fN]	
targets	[f1 ... fN]	
candidate_inputs	[f1 ... fN]	
forced_inputs	[f1 ... fN]	
use_timestamp	Timestamp	
	Period	

表 149. tcmnode プロパティ (続き)

tcmnode プロパティ	値	プロパティの説明
input_interval	None Unknown Year Quarter Month Week Day Hour Hour_nonperiod Minute Minute_nonperiod Second Second_nonperiod	
period_field	<i>string</i>	
period_start_value	<i>integer</i>	
num_days_per_week	<i>integer</i>	
start_day_of_week	Sunday Monday Tuesday Wednesday Thursday Friday Saturday	
num_hours_per_day	<i>integer</i>	
start_hour_of_day	<i>integer</i>	
timestamp_increments	<i>integer</i>	

表 149. *tcmnode* プロパティ (続き)

tcmnode プロパティ	値	プロパティの説明
cyclic_increments	<i>integer</i>	
cyclic_periods	リスト	
output_interval	None Year Quarter Month Week Day Hour Minute Second	
is_same_interval	Same Notsame	
cross_hour	<i>Boolean</i>	
aggregate_and_distribute	リスト	
aggregate_default	Mean Sum Mode Min Max	
distribute_default	Mean Sum	

表 149. tcmnode プロパティ (続き)

tcmnode プロパティ	値	プロパティの説明
group_default	Mean Sum Mode Min Max	
missing_imput	Linear_interp Series_mean K_mean K_meridian Linear_trend None	
k_mean_param	integer	
k_median_param	integer	
missing_value_threshold	integer	
conf_level	integer	
max_num_predictor	integer	
max_lag	integer	
epsilon	number	
threshold	integer	
is_re_est	Boolean	
num_targets	integer	
percent_targets	integer	
fields_display	リスト	
series_display	リスト	
network_graph_for_target	Boolean	
sign_level_for_target	number	
fit_and_outlier_for_target	Boolean	
sum_and_para_for_target	Boolean	
impact_diag_for_target	Boolean	

表 149. tcmnode プロパティ (続き)

tcmnode プロパティ	値	プロパティの説明
impact_diag_type_for_target	Effect Cause Both	
impact_diag_level_for_target	integer	
series_plot_for_target	Boolean	
res_plot_for_target	Boolean	
top_input_for_target	Boolean	
forecast_table_for_target	Boolean	
same_as_for_target	Boolean	
network_graph_for_series	Boolean	
sign_level_for_series	number	
fit_and_outlier_for_series	Boolean	
sum_and_para_for_series	Boolean	
impact_diagram_for_series	Boolean	
impact_diagram_type_for_series	Effect Cause Both	
impact_diagram_level_for_series	integer	
series_plot_for_series	Boolean	
residual_plot_for_series	Boolean	
forecast_table_for_series	Boolean	
outlier_root_cause_analysis	Boolean	
causal_levels	integer	
outlier_table	Interactive Pivot Both	
rmisp_error	Boolean	

表 149. tcmnode プロパティ (続き)

tcmnode プロパティ	値	プロパティの説明
bic	Boolean	
r_square	Boolean	
outliers_over_time	Boolean	
series_transormation	Boolean	
use_estimation_period	Boolean	
estimation_period	Times Observation	
observations	リスト	
observations_type	Latest Earliest	
observations_num	integer	
observations_exclude	integer	
extend_records_into_future	Boolean	
forecastperiods	integer	
max_num_distinct_values	integer	
display_targets	FIXEDNUMBER PERCENTAGE	
goodness_fit_measure	ROOTMEAN BIC RSQUARE	
top_input_for_series	Boolean	
aic	Boolean	
rmse	Boolean	

ts プロパティ



時系列ノードは、時系列データから指数平滑法、1 変量の自己回帰型統合移動平均法 (ARIMA)、および多変量 ARIMA (または伝達関数) モデルを推測し、将来のパフォーマンスの予測を作成します。この時系列ノードは、スpos モデラー バージョン 18 で廃止された以前の時系列ノードと類似しています。ただし、この新しい時系列ノードは、IBMSPSS Analytic Server の機能を活用してビッグデータを処理するよう設計されており、結果モデルは スpos モデラー バージョン 17 で追加された出力ビューアーに表示されます。

表 150. ts プロパティ

ts プロパティ	値	プロパティの説明
targets	フィールド	時系列ノードは、オプションで 1 つ以上の入力フィールドを予測値として使用しながら、1 つ以上の対象フィールドを予測します。度数フィールドおよび重みフィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
candidate_inputs	[field1 ... fieldN]	モデルで使用される入力または予測変数フィールド。
use_period	フラグ	
date_time_field	フィールド	
input_interval	None Unknown Year Quarter Month Week Day Hour Hour_nonperiod Minute Minute_nonperiod Second Second_nonperiod	
period_field	フィールド	
period_start_value	integer	
num_days_per_week	integer	

表 150. ts プロパティ (続き)		
ts プロパティ	値	プロパティの説明
start_day_of_week	Sunday Monday Tuesday Wednesday Thursday Friday Saturday	
num_hours_per_day	<i>integer</i>	
start_hour_of_day	<i>integer</i>	
timestamp_increments	<i>integer</i>	
cyclic_increments	<i>integer</i>	
cyclic_periods	リスト	
output_interval	None Year Quarter Month Week Day Hour Minute Second	
is_same_interval	フラグ	
cross_hour	フラグ	
aggregate_and_distribute	リスト	

表 150. ts プロパティ (続き)

ts プロパティ	値	プロパティの説明
aggregate_default	Mean	
	Sum	
	Mode	
	Min	
	Max	
distribute_default	Mean	
	Sum	
group_default	Mean	
	Sum	
	Mode	
	Min	
	Max	
missing_imput	Linear_interp	
	Series_mean	
	K_mean	
	K_median	
	Linear_trend	
k_span_points	integer	
use_estimation_period	フラグ	
estimation_period	Observations	
	Times	
date_estimation	リスト	date_time_field を使用する場合にのみ使用可能です
period_estimation	リスト	use_period を使用する場合にのみ使用可能です
observations_type	Latest	
	Earliest	

表 150. ts プロパティ (続き)		
ts プロパティ	値	プロパティの説明
observations_num	<i>integer</i>	
observations_exclude	<i>integer</i>	
method	ExpertModeler Exsmooth Arima	
expert_modeler_method	ExpertModeler Exsmooth Arima	
consider_seasonal	フラグ	
detect_outliers	フラグ	
expert_outlier_additive	フラグ	
expert_outlier_level_shift	フラグ	
expert_outlier_innovational	フラグ	
expert_outlier_level_shift	フラグ	
expert_outlier_transient	フラグ	
expert_outlier_seasonal_additive	フラグ	
expert_outlier_local_trend	フラグ	
expert_outlier_additive_patch	フラグ	
consider_newesmodels	フラグ	

表 150. ts プロパティ (続き)

ts プロパティ	値	プロパティの説明
exsmooth_model_type	Simple HoltsLinearTrend BrownsLinearTrend DampedTrend SimpleSeasonal WintersAdditive WintersMultiplicative DampedTrendAdditive DampedTrendMultiplicative MultiplicativeTrendAdditive MultiplicativeSeasonal MultiplicativeTrendMultiplicative MultiplicativeTrend	指数平滑法を指定します。デフォルトは Simple です。

表 150. ts プロパティ (続き)

ts プロパティ	値	プロパティの説明
futureValue_type_method	Compute specify	Compute を使用すると、各予測の予測期間に対し、将来の値が計算されます。 予測変数ごとに、関数のリスト (空白、最近のポイントの平均、最新の値) から選択することも、specify を使用して手動で値を入力することもできます。個々のフィールドとプロパティを指定するには、extend_metric_values プロパティを使用します。以下に例を示します。 <pre>set :ts.futureValue_type_method="specify" set :ts.extend_metric_values=[{'Market_1', 'USER_SPECIFY', [1,2,3]}, {'Market_2', 'MOST_RECENT_VALUE', ''}, {'Market_3', 'RECENT_POINTS_MEAN', ''}]</pre>
exsmooth_transformation_type	None SquareRoot NaturalLog	
arma.p	integer	
arma.d	integer	
arma.q	integer	
arma.sp	integer	
arma.sd	integer	
arma.sq	integer	
arma_transformation_type	None SquareRoot NaturalLog	
arma_include_constant	フラグ	
tf_arma.p.fieldname	integer	伝達関数用。
tf_arma.d.fieldname	integer	伝達関数用。

表 150. ts プロパティ (続き)		
ts プロパティ	値	プロパティの説明
tf_arma.q.fieldname	integer	伝達関数用。
tf_arma.sp.fieldname	integer	伝達関数用。
tf_arma.sd.fieldname	integer	伝達関数用。
tf_arma.sq.fieldname	integer	伝達関数用。
tf_arma.delay.fieldname	integer	伝達関数用。
tf_arma.transformation_type. fieldname	None SquareRoot NaturalLog	伝達関数用。
arma_detect_outliers	フラグ	
arma_outlier_additive	フラグ	
arma_outlier_level_shift	フラグ	
arma_outlier_innovational	フラグ	
arma_outlier_transient	フラグ	
arma_outlier_seasonal_additive	フラグ	
arma_outlier_local_trend	フラグ	
arma_outlier_additive_patch	フラグ	
max_lags	integer	
cal_PI	フラグ	
conf_limit_pct	real	
events	フィールド	
continue	フラグ	
scoring_model_only	フラグ	多く (1 万単位) の時系列 のモデルに使用します。
forecastperiods	integer	
extend_records_into_future	フラグ	
extend_metric_values	フィールド	予測のための将来の値を 提供できるようにしま す。
conf_limits	フラグ	
noise_res	フラグ	

表 150. ts プロパティ (続き)		
ts プロパティ	値	プロパティの説明
max_models_output	integer	出力に表示するモデルの数を制御します。デフォルトは 10 です。構築されたモデルの総数がこの値を超える場合、モデルは出力に表示されません。それでも、モデルはスコアリングに引き続き使用できます。

timeseriesnode プロパティ (廃止)



注: この元の時系列ノードは スポス モデラー のバージョン 18 で廃止され、新しい時系列ノードに置き換えられました。この新しいノードは、IBMSPPS Analytic Server の機能を活用し、ビッグデータを処理するように設計されています。

時系列ノードは、時系列データから指数平滑法、1 変量の自己回帰型統合移動平均法 (ARIMA)、および多変量 ARIMA (または伝達関数) モデルを推測し、将来のパフォーマンスの予測を作成します。時系列ノードは、時間区分ノードによって常に先行される必要があります。

例

```
node = stream.create("timeseries", "My node")
node.setPropertyValue("method", "Exsmooth")
node.setPropertyValue("exsmooth_model_type", "HoltsLinearTrend")
node.setPropertyValue("exsmooth_transformation_type", "None")
```

表 151. timeseriesnode プロパティ		
timeseriesnode プロパティ	値	プロパティの説明
targets	フィールド	時系列ノードは、オプションで 1 つ以上の入力フィールドを予測値として使用しながら、1 つ以上の対象フィールドを予測します。度数フィールドおよび重みフィールドは使用しません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
continue	フラグ	
method	ExpertModeler Exsmooth Arima Reuse	

表 151. *timeseriesnode* プロパティ (続き)

timeseriesnode プロパティ	値	プロパティの説明
expert_modeler_method	フラグ	
consider_seasonal	フラグ	
detect_outliers	フラグ	
expert_outlier_additive	フラグ	
expert_outlier_level_shift	フラグ	
expert_outlier_innovational	フラグ	
expert_outlier_level_shift	フラグ	
expert_outlier_transient	フラグ	
expert_outlier_seasonal_additive	フラグ	
expert_outlier_local_trend	フラグ	
expert_outlier_additive_patch	フラグ	
exsmooth_model_type	Simple HoltsLinearTrend BrownsLinearTrend DampedTrend SimpleSeasonal WintersAdditive WintersMultiplicative	
exsmooth_transformation_type	None SquareRoot NaturalLog	
arma_p	<i>integer</i>	
arma_d	<i>integer</i>	
arma_q	<i>integer</i>	
arma_sp	<i>integer</i>	
arma_sd	<i>integer</i>	
arma_sq	<i>integer</i>	

表 151. *timeseriesnode* プロパティ (続き)

timeseriesnode プロパティ	値	プロパティの説明
arma_transformation_type	None SquareRoot NaturalLog	
arma_include_constant	フラグ	
tf_arma_p.fieldname	integer	伝達関数用。
tf_arma_d.fieldname	integer	伝達関数用。
tf_arma_q.fieldname	integer	伝達関数用。
tf_arma_sp.fieldname	integer	伝達関数用。
tf_arma_sd.fieldname	integer	伝達関数用。
tf_arma_sq.fieldname	integer	伝達関数用。
tf_arma_delay.fieldname	integer	伝達関数用。
tf_arma_transformation_type. fieldname	None SquareRoot NaturalLog	伝達関数用。
arma_detect_outlier_mode	None Automatic	
arma_outlier_additive	フラグ	
arma_outlier_level_shift	フラグ	
arma_outlier_innovational	フラグ	
arma_outlier_transient	フラグ	
arma_outlier_seasonal_additive	フラグ	
arma_outlier_local_trend	フラグ	
arma_outlier_additive_patch	フラグ	
conf_limit_pct	real	
max_lags	integer	
events	フィールド	
scoring_model_only	フラグ	多く (1 万単位) の時系列 のモデルに使用します。

treeas プロパティ



Tree-AS ノードは既存の CHAID ノードに似ていますが、Tree-AS ノードはビッグデータを処理して 1 つのツリーを作成することを目的に設計されており、結果モデルがスポンサー モデラー バージョン 17 で追加された出力ビューアーに表示されます。このノードは、カイ 2 乗統計量 (CHAID) を使用して最適な分割を特定することで、ディビジョン・ツリーを生成します。CHAID をこのように使用することで、非 2 分岐ツリーを生成できます。これは、3 個以上のブランチを持つ分岐が存在することを意味します。対象フィールドおよび入力フィールドは、数値範囲 (連続型) またはカテゴリーとなります。Exhaustive CHAID は CHAID の修正版で、可能性のある分割すべてを調べることで、よりよい結果を得られますが、計算時間も長くなります。

表 152. treeas プロパティ

treeas プロパティ	値	プロパティの説明
target	フィールド	Tree-AS ノードでは、CHAID モデルには単一の対象フィールドおよび 1 つ以上の入力フィールドが必要になります。度数フィールドも指定できます。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
method	chaid exhaustive_chaid	
max_depth	integer	最大ツリー深度 (0 から 20)。デフォルト値は 5 です。
num_bins	integer	データが連続型入力で構成される場合にのみ使用されます。入力に対して使用する等しいフリクエンシ ビンの数を設定します。オプションは 2、4、5、10、20、25、50、または 100 です。
record_threshold	integer	モデルでツリーを作成するときに、p 値の使用から効果サイズの使用に切り替えるレコード数。デフォルトは 1,000,000 です。増減は 10,000 の単位で行います。
split_alpha	number	分割の有意水準。この値は 0.01 から 0.99 までです。
merge_alpha	number	結合の有意水準。この値は 0.01 から 0.99 までです。
bonferroni_adjustment	フラグ	Bonferroni メソッドを使用して有意確率値を調整。
effect_size_threshold_continuous	number	連続型対象を使用する際にノードの分割およびカテゴリの結合を行う効果サイズしきい値を設定します。この値は 0.01 から 0.99 までです。
effect_size_threshold_categorical	number	カテゴリ型対象を使用する際にノードの分割およびカテゴリの結合を行う効果サイズしきい値を設定します。この値は 0.01 から 0.99 までです。

表 152. *treeas* プロパティ (続き)

treeas プロパティ	値	プロパティの説明
split_merged_categories	フラグ	マージしたカテゴリーの再分割を許可。
grouping_sig_level	<i>number</i>	ノードグループの形成方法または例外ノードの識別方法を決定するために使用されます。
chi_square	pearson likelihood_ratio	カイ 2 乗統計の計算に使用される方法 (Pearson または尤度比)
minimum_record_use	use_percentage use_absolute	
min_parent_records_pc	<i>number</i>	デフォルト値は 2 です。最小は 1、最大は 100、インクリメントは 1 です。親ブランチの値は子ブランチより高くなければなりません。
min_child_records_pc	<i>number</i>	デフォルト値は 1 です。最小は 1、最大は 100、インクリメントは 1 です。
min_parent_records_abs	<i>number</i>	デフォルト値は 100 です。最小は 1、最大は 100、インクリメントは 1 です。親ブランチの値は子ブランチより高くなければなりません。
min_child_records_abs	<i>number</i>	デフォルト値は 50 です。最小は 1、最大は 100、インクリメントは 1 です。
epsilon	<i>number</i>	期待されるセル度数の最小変化量。
max_iterations	<i>number</i>	収束のための最大反復回数。
use_costs	フラグ	
costs	構造化	構造化プロパティ。形式は、実際の値、予測された値、およびコスト (予測が正しくない場合) の 3 つの値のリストです。以下に例を示します。 <code>tree.setPropertyValue("costs", [["drugA", "drugB", 3.0], ["drugX", "drugY", 4.0]])</code>
default_cost_increase	none linear square custom	注: 順序型対象に対してのみ有効です。 コスト行列にデフォルト値を設定します。
calculate_conf	フラグ	

表 152. *treeas* プロパティ (続き)

treeas プロパティ	値	プロパティの説明
display_rule_id	フラグ	フィールドが 1 つスコアリング出力に追加されますが、これは各レコードを割り当てるターミナル・ノードに ID を示すためのものです。

twostepnode プロパティ



TwoStep ノードで、2 段階のクラスター化手法が使用されます。最初のステップでは、データを 1 度通過させて、未処理の入力データを管理可能な一連のサブクラスターに圧縮します。2 番目のステップでは、階層クラスター化手法を使用して、サブクラスターをより大きなクラスターに結合させていきます。TwoStep には、学習データに最適なクラスター数を自動的に推定するという利点があります。また、フィールド・タイプの混在や大規模データ・セットも効率よく処理できます。

例

```
node = stream.create("twostep", "My node")
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("inputs", ["Age", "K", "Na", "BP"])
node.setPropertyValue("partition", "Test")
node.setPropertyValue("use_model_name", False)
node.setPropertyValue("model_name", "TwoStep_Drug")
node.setPropertyValue("use_partitioned_data", True)
node.setPropertyValue("exclude_outliers", True)
node.setPropertyValue("cluster_label", "String")
node.setPropertyValue("label_prefix", "TwoStep_")
node.setPropertyValue("cluster_num_auto", False)
node.setPropertyValue("max_num_clusters", 9)
node.setPropertyValue("min_num_clusters", 3)
node.setPropertyValue("num_clusters", 7)
```

表 153. *twostepnode* プロパティ

twostepnode プロパティ	値	プロパティの説明
inputs	[<i>field1 ... fieldN</i>]	TwoStep モデルは対象フィールドでなく、入力フィールドのリストを使用します。重みフィールドおよび度数フィールドは認識されません。詳しくは、トピック 217 ページの『一般的なモデル作成ノードのプロパティ』を参照してください。
standardize	フラグ	
exclude_outliers	フラグ	
percentage	<i>number</i>	
cluster_num_auto	フラグ	
min_num_clusters	<i>number</i>	
max_num_clusters	<i>number</i>	
num_clusters	<i>number</i>	

表 153. *twostepnode* プロパティ (続き)

twostepnode プロパティ	値	プロパティの説明
cluster_label	String	
	Number	
label_prefix	string	
distance_measure	Euclidean	
	Loglikelihood	
clustering_criterion	AIC	
	BIC	

twostepAS のプロパティ



TwoStep クラスターは、通常は明確ではない、データ・セット内での自然なグループ化 (またはクラスター) を明確にすることを目的として設計された探索ツールです。この手続きで使用するアルゴリズムには、従来のクラスタリング手法とは異なる以下の優れた特徴があります (カテゴリー変数および連続変数の処理、クラスター数の自動選択、スケーラビリティなど)。

表 154. *twostepAS* のプロパティ

twostepAS のプロパティ	値	プロパティの説明
inputs	[$f_1 \dots f_N$]	TwoStepAS モデルは入力フィールドのリストを使用しますが、対象フィールドは使用しません。重みフィールドおよび度数フィールドは認識されません。
use_predefined_roles	ブール値	デフォルト=True
use_custom_field_assignments	ブール値	デフォルト=False
cluster_num_auto	ブール値	デフォルト=True
min_num_clusters	整数	Default=2
max_num_clusters	整数	Default=15
num_clusters	整数	Default=5
clustering_criterion	AIC	
	BIC	

表 154. twostepAS のプロパティ (続き)

twostepAS のプロパティ	値	プロパティの説明
automatic_clustering_method	use_clustering_criterion_setting Distance_jump Minimum Maximum	
feature_importance_method	use_clustering_criterion_setting effect_size	
use_random_seed	ブール値	
random_seed	整数	
distance_measure	Euclidean Loglikelihood	
include_outlier_clusters	ブール値	デフォルト=True
num_cases_in_feature_tree_leaf_is_less_than	整数	Default=10
top_perc_outliers	整数	Default=5
initial_dist_change_threshold	整数	Default=0
leaf_node_maximum_branches	整数	Default=8
non_leaf_node_maximum_branches	整数	Default=8
max_tree_depth	整数	Default=3
adjustment_weight_on_measurement_level	整数	Default=6
memory_allocation_mb	number	Default=512
delayed_split	ブール値	デフォルト=True
fields_to_standardize	[f1 ... fN]	
adaptive_feature_selection	ブール値	デフォルト=True
featureMisPercent	整数	Default=70
coefRange	number	Default=0.05
percCasesSingleCategory	整数	Default=95
numCases	整数	Default=24
include_model_specifications	ブール値	デフォルト=True
include_record_summary	ブール値	デフォルト=True
include_field_transformations	ブール値	デフォルト=True

表 154. <i>twostepAS</i> のプロパティ (続き)		
twostepAS のプロパティ	値	プロパティの説明
excluded_inputs	ブール値	デフォルト=True
evaluate_model_quality	ブール値	デフォルト=True
show_feature_importance_bar chart	ブール値	デフォルト=True
show_feature_importance_ word_cloud	ブール値	デフォルト=True
show_outlier_clusters interactive_table_and_chart	ブール値	デフォルト=True
show_outlier_clusters_pivot_table	ブール値	デフォルト=True
across_cluster_feature_importance	ブール値	デフォルト=True
across_cluster_profiles_pivot_table	ブール値	デフォルト=True
withinprofiles	ブール値	デフォルト=True
cluster_distances	ブール値	デフォルト=True
cluster_label	String	
	Number	
label_prefix	String	

第 14 章 モデル・ナゲット・ノードのプロパティ

モデル・ナゲット・ノードは、他のノードと同じ共通のプロパティを共有しています。詳しくは、トピック 76 ページの『共通のノード・プロパティ』を参照してください。

applyanomalydetectionnode プロパティ

異常値検査モデル作成ノードを使用して、異常値検査モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyanomalydetectionnode* です。モデル作成ノード自体のスクリプトの詳細は、218 ページの『anomalydetectionnode プロパティ』を参照してください。

表 155. <i>applyanomalydetectionnode</i> プロパティ		
applyanomalydetectionnode プロパティ	値	プロパティの説明
anomaly_score_method	FlagAndScore FlagOnly ScoreOnly	スコアリング用に、作成される出力を決めます。
num_fields	integer	報告するフィールド数。
discard_records	フラグ	レコードが出力から廃棄されるかどうかを示します。
discard_anomalous_records	フラグ	異常なレコードを廃棄するか、または異常でないレコードを廃棄するかの標識。デフォルトは、異常でないレコードが廃棄されることを示す off です。それに対し、 on の場合は、異常なレコードが廃棄されます。このプロパティは、 discard_records が有効な場合にだけ、有効になります。

applyapriorinode プロパティ

Apriori モデル作成ノードを使用して、Apriori モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyapriorinode* です。モデル作成ノード自体のスクリプトの詳細は、219 ページの『apriorinode プロパティ』を参照してください。

表 156. <i>applyapriorinode</i> プロパティ		
applyapriorinode プロパティ	値	プロパティの説明
max_predictions	数値 (整数)	
ignore_unmatached	フラグ	
allow_repeats	フラグ	
check_basket	NoPredictions Predictions NoCheck	

表 156. *applyapriorinode* プロパティ (続き)

applyapriorinode プロパティ	値	プロパティの説明
criterion	Confidence Support RuleSupport Lift Deployability	

applyassociationrulesnode プロパティ

アソシエーションルールモデル作成ノードを使用して、アソシエーションルールモデルナゲットを作成することができます。このモデルナゲットのスクリプト名は *applyassociationrulesnode* です。モデル作成ノード自体をスクリプト化する方法については、[221 ページの『associationrulesnode プロパティ』](#)を参照してください。

表 157. *applyassociationrulesnode* プロパティ

applyassociationrulesnode プロパティ	データ型	プロパティの説明
max_predictions	integer	スコアに対する各入力に適用できるルールの最大数。
criterion	Confidence Rulesupport Lift Conditionsupport Deployability	ルールの強度を判断するための尺度を選択します。
allow_repeats	Boolean	同じ予測を持つルールをスコア内に含めるかどうかを決定します。
check_input	NoPredictions Predictions NoCheck	

applyautoclassifiernode プロパティ

自動分類モデル作成ノードを使用して、自動分類モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyautoclassifiernode* です。モデル作成ノードのスクリプト化の詳細は、[224 ページの『autoclassifiernode プロパティ』](#)を参照してください。

表 158. applyautoclassifiernode プロパティ

applyautoclassifiernode プロパティ	値	プロパティの説明
flag_ensemble_method	Voting EvaluationWeightedVoting ConfidenceWeightedVoting RawPropensityWeightedVoting HighestConfidence AverageRawPropensity	アンサンブル・スコアを決定するために使用する方法を指定します。この設定は、選択された対象がフラグ型フィールドである場合にのみ適用されます。
flag_evaluation_selection	Accuracy AUC_ROC	このオプションは、評価-重み付き票決にどの評価測定を選択するかを決定するためのフラグ型対象専用です。
filter_individual_model_output	フラグ	個々のモデルのスコアリング結果を抑制するかどうかを指定します。
is_ensemble_update	フラグ	継続的自動機械学習モードを有効にします。これにより、既存の自動モデルを置き換えるのではなく、新しいコンポーネント・モデルが既存の自動モデルに追加され、新たに使用可能になったデータを使用して既存のコンポーネント・モデルの尺度が再評価されます。
is_auto_ensemble_weights_reevaluation	フラグ	自動モデルの重みの再評価を有効にします。
use_accumulated_factor	フラグ	累積係数を使用して、累積尺度を計算します。
accumulated_factor	数値 (double)	最大値は 0.99 で、最小値は 0.85 です。
use_accumulated_reducing	フラグ	モデルの更新時に累積制限に基づいてモデルの縮小を実行します。
accumulated_reducing_limit	数値 (double)	最大値は 0.7 で、最小値は 0.1 です。
use_accumulated_weighted_evaluation	フラグ	アンサンブル法として評価-重み付き票決法を選択した場合に票決で累積評価測定が使用されます。

表 158. *applyautoclassifiernode* プロパティ (続き)

applyautoclassifiernode プロパティ	値	プロパティの説明
flag_voting_tie_selection	Random HighestConfidence RawPropensity	票決方法が選択された場合、可否同数の解決方法を指定します。この設定は、選択された対象がフラグ型フィールドである場合にのみ適用されます。
set_ensemble_method	Voting EvaluationWeightedVoting ConfidenceWeightedVoting HighestConfidence	アンサンブル・スコアを決定するために使用する方法を指定します。この設定は、選択された対象がセット型フィールドである場合にのみ適用されます。
set_voting_tie_selection	Random HighestConfidence	票決方法が選択された場合、可否同数の解決方法を指定します。この設定は、選択された対象が名義型フィールドである場合にのみ適用されます。

applyautoclusternode プロパティ

自動クラスタリング・モデル作成ノードを使用して、自動クラスタリング・モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyautoclusternode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体のスクリプトの詳細は、[226 ページの『autoclusternode プロパティ』](#)を参照してください。

applyautonumericnode プロパティ

自動数値モデル作成ノードを使用して、自動数値モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyautonumericnode* です。モデル作成ノードのスクリプト化の詳細は、[228 ページの『autonumericnode プロパティ』](#)を参照してください。

表 159. *applyautonumericnode* プロパティ

applyautonumericnode プロパティ	値	プロパティの説明
calculate_standard_error	フラグ	

applybayesnetnode プロパティ

Bayesian network (ベイズ) モデル作成ノードを使用して、Bayesian network (ベイズ) モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applybayesnetnode* です。モデル作成ノード自体のスクリプトの詳細は、[230 ページの『bayesnetnode プロパティ』](#)を参照してください。

表 160. *applybayesnetnode* プロパティ

applybayesnetnode プロパティ	値	プロパティの説明
all_probabilities	フラグ	

表 160. *applybayesnetnode* プロパティ (続き)

applybayesnetnode プロパティ	値	プロパティの説明
raw_propensity	フラグ	
adjusted_propensity	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applyc50node プロパティ

C5.0 モデル作成ノードを使用して、C5.0 モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyc50node* です。モデル作成ノード自体のスクリプトの詳細は、[232 ページの『c50node プロパティ』](#)を参照してください。

表 161. *applyc50node* プロパティ

applyc50node プロパティ	値	プロパティの説明
sql_generate	udf Never NoMissingValues	ルールセット実行時の SQL 生成オプションの設定に使用します。デフォルト値は udf です。
calculate_conf	フラグ	SQL 生成が有効になっている場合に利用できます。このプロパティには、生成されたツリー中の確信度計算が含まれています。
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applycarmanode プロパティ

CARMA モデル作成ノードを使用して、CARMA モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applycarmanode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体のスクリプトの詳細は、[234 ページの『carmanode プロパティ』](#)を参照してください。

applycartnode プロパティ

C&R Tree モデル作成ノードを使用して、C&R Tree モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applycartnode* です。モデル作成ノード自体のスクリプトの詳細は、[235 ページの『cartnode プロパティ』](#)を参照してください。

表 162. *applycartnode* プロパティ

applycartnode プロパティ	値	プロパティの説明
enable_sql_generation	Never MissingValues NoMissingValues	ルールセット実行時の SQL 生成オプションの設定に使用します。
calculate_conf	フラグ	SQL 生成が有効になっている場合に利用できます。このプロパティには、生成されたツリー中の確信度計算が含まれています。
display_rule_id	フラグ	フィールドが 1 つスコアリング出力に追加されますが、これは各レコードを割り当てるターミナル・ノードに ID を示すためのものです。
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applychaidnode プロパティ

CHAID モデル作成ノードを使用して、CHAID モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applychaidnode* です。モデル作成ノード自体のスクリプトの詳細は、[238 ページ](#)の『*chaidnode* プロパティ』を参照してください。

表 163. *applychaidnode* プロパティ

applychaidnode Properties	値	プロパティの説明
enable_sql_generation	Never MissingValues	ルールセット実行時の SQL 生成オプションの設定に使用します。
calculate_conf	フラグ	
display_rule_id	フラグ	フィールドが 1 つスコアリング出力に追加されますが、これは各レコードを割り当てるターミナル・ノードに ID を示すためのものです。
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applycoxregnode プロパティ

Cox モデル作成ノードを使用して、Cox モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applycoxregnode* です。モデル作成ノード自体のスクリプトの詳細は、[240 ページ](#)の『*coxregnode* プロパティ』を参照してください。

表 164. *applycoxregnode* プロパティ

applycoxregnode プロパティ	値	プロパティの説明
future_time_as	Intervals Fields	
time_interval	<i>number</i>	
num_future_times	<i>integer</i>	
time_field	フィールド	
past_survival_time	フィールド	
all_probabilities	フラグ	
cumulative_hazard	フラグ	

applydecisionlistnode プロパティ

ディシジョン・リスト・モデル作成ノードを使用して、ディシジョン・リスト・モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applydecisionlistnode* です。モデル作成ノード自体のスクリプトの詳細は、[243 ページの『decisionlistnode プロパティ』](#)を参照してください。

表 165. *applydecisionlistnode* プロパティ

applydecisionlistnode プロパティ	値	プロパティの説明
enable_sql_generation	フラグ	真に設定したときは、ディシジョン・リスト・モデルが SQL ヘプッシュバックされるように IBM スポス モデラー が試行します。
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applydiscriminantnode プロパティ

判別分析モデル作成ノードを使用して、判別分析モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applydiscriminantnode* です。モデル作成ノード自体のスクリプトの詳細は、[244 ページの『discriminantnode プロパティ』](#)を参照してください。

表 166. *applydiscriminantnode* プロパティ

applydiscriminantnode プロパティ	値	プロパティの説明
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	



拡張モデル・ノードを使用して、拡張モデル・ナゲットを生成することができます。このモデルナゲットのスクリプト名は *applyextension* です。モデル作成ノード自体をスクリプト化する方法については、[246 ページの『extensionmodelnode プロパティ』](#)を参照してください。

Python for Spark の例

```
##### script example for Python for Spark
applyModel = stream.findByType("extension_apply", None)

score_script = """
import json
import spss.pyspark.runtime
from pyspark.mllib.regression import LabeledPoint
from pyspark.mllib.linalg import DenseVector
from pyspark.mllib.tree import DecisionTreeModel
from pyspark.sql.types import StringType, StructField

cxt = spss.pyspark.runtime.getContext()

if cxt.isComputeDataModelOnly():
    _schema = cxt.getSparkInputSchema()
    _schema.fields.append(StructField("Prediction", StringType(), nullable=True))
    cxt.setSparkOutputSchema(_schema)
else:
    df = cxt.getSparkInputData()

    _modelPath = cxt.getModelContentToPath("TreeModel")
    metadata = json.loads(cxt.getModelContentToString("model.dm"))

    schema = df.dtypes[:]
    target = "Drug"
    predictors = ["Age", "BP", "Sex", "Cholesterol", "Na", "K"]

    lookup = {}
    for i in range(0, len(schema)):
        lookup[schema[i][0]] = i

    def row2LabeledPoint(dm, lookup, target, predictors, row):
        target_index = lookup[target]
        tval = dm[target_index].index(row[target_index])
        pvals = []
        for predictor in predictors:
            predictor_index = lookup[predictor]
            if isinstance(dm[predictor_index], list):
                pval = row[predictor_index] in dm[predictor_index] and
                dm[predictor_index].index(row[predictor_index]) or -1
            else:
                pval = row[predictor_index]
            pvals.append(pval)
        return LabeledPoint(tval, DenseVector(pvals))

    # convert dataframe to an RDD containing LabeledPoint
    lps = df.rdd.map(lambda row: row2LabeledPoint(metadata, lookup, target, predictors, row))
    treeModel = DecisionTreeModel.load(cxt.getSparkContext(), _modelPath);
    # score the model, produces an RDD containing just double values
    predictions = treeModel.predict(lps.map(lambda lp: lp.features))

    def addPrediction(x, dm, lookup, target):
        result = []
        for _idx in range(0, len(x[0])):
            result.append(x[0][_idx])
            result.append(dm[lookup[target]][int(x[1])])
        return result

    _schema = cxt.getSparkInputSchema()
    _schema.fields.append(StructField("Prediction", StringType(), nullable=True))
    rdd2 = df.rdd.zip(predictions).map(lambda x: addPrediction(x, metadata, lookup, target))
    outDF = cxt.getSparkSQLContext().createDataFrame(rdd2, _schema)

    cxt.setSparkOutputData(outDF)
```

```
"""
applyModel.setPropertyValue("python_syntax", score_script)
```

R の例

```
#### script example for R
applyModel.setPropertyValue("r_syntax", """
result<-predict(modelerModel,newdata=modelerData)
modelerData<-cbind(modelerData,result)
var1<-c(fieldName="NaPrediction",fieldLabel="",fieldStorage="real",fieldMeasure="",
fieldFormat="",fieldRole="")
modelerDataModel<-data.frame(modelerDataModel,var1)""")
```

表 167. *applyextension* プロパティ

applyextension プロパティ	値	プロパティの説明
r_syntax	string	モデル・スコアリング用の R スクリプト・シンタックス。
python_syntax	string	モデル・スコアリング用の Python スクリプト・シンタックス。
use_batch_size	フラグ	バッチ処理を使用可能にします。
batch_size	integer	各バッチに含めるデータレコードの数を指定します。
convert_flags	StringsAndDoubles LogicalValues	フラグ型フィールドを変換するためのオプション。
convert_missing	フラグ	欠損値を R の NA 値に変換するオプションです。
convert_datetime	フラグ	日付形式または日付/時刻形式の変数を R の日付/時刻形式に変換するためのオプション。
convert_datetime_class	POSIXct POSIXlt	日付形式または日付/時刻形式の変数のうち、どの形式の変数を変換するかを指定するためのオプション。

applyfactornode プロパティ

因子分析モデル作成ノードを使用して、因子分析モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyfactornode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体のスクリプトの詳細は、249 ページの『*factornode* プロパティ』を参照してください。

applyfeatureselectionnode プロパティ

特微量選択モデル作成ノードを使用して、特微量選択モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyfeatureselectionnode* です。モデル作成ノード自体のスクリプトの詳細は、251 ページの『*featureselectionnode* プロパティ』を参照してください。

表 168. *applyfeatureselectionnode* プロパティ

applyfeatureselectionnode プロパティ	値	プロパティの説明
selected_ranked_fields		モデル・ブラウザー内で検査されるランク付きのフィールドを指定します。
selected_screened_fields		モデル・ブラウザー内で検査されるスクリーニングされたフィールドを指定します。

applygeneralizedlinearnode プロパティ

一般化線型 (genlin) モデル作成ノードを使用して、一般化線型モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applygeneralizedlinearnode* です。モデル作成ノード自体のスクリプトの詳細は、[253 ページの『genlinnode プロパティ』](#)を参照してください。

表 169. *applygeneralizedlinearnode* プロパティ

applygeneralizedlinearnode プロパティ	値	プロパティの説明
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applyglmnode プロパティ

GLMM モデル作成ノードを使用して、GLMM モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyglmnode* です。モデル作成ノード自体のスクリプトの詳細は、[258 ページの『glmnode プロパティ』](#)を参照してください。

表 170. *applyglmnode* プロパティ

applyglmnode プロパティ	値	プロパティの説明
confidence	onProbability onIncrease	スコアリングの確信度を計算する基準 (最も高い予測確率、または最も高い予測確率と 2 番目に高い予測確率との差)。
score_category_probabilities	フラグ	True に設定された場合、カテゴリ対象の予測確率を生成します。カテゴリごとにフィールドが作成されます。デフォルトは False です。
max_categories	integer	確率を予測するカテゴリの最大数です。 <i>score_category_probabilities</i> が True の場合にのみ使用されます。

表 170. *applyglmnode* プロパティ (続き)

applyglmnode プロパティ	値	プロパティの説明
score_propensity	フラグ	True に設定された場合、フラグ型対象を含むモデルに対して、未調整傾向スコア (「true」の結果の確率) を生成します。データ区分が有効な場合、テスト・データ区分に基づいて、調整済み傾向スコアも生成します。デフォルトは False です。
enable_sql_generation	udf native	ストリーム実行中の SQL 生成オプションを設定するために使用します。データベースにプッシュバックして SPSS® Modeler Server Scoring Adapter でスコアリングするか (スコアリングアダプターがインストール済みのデータベースに接続している場合)、SPSS Modeler 内でスコアリングするかを選択できます。 デフォルト値は udf です。

applygle プロパティ

GLE モデル作成ノードを使用して、GLE モデルナゲットを生成できます。このモデルナゲットのスクリプト名は *applygle* です。モデル作成ノード自体をスクリプト化する方法については、[263 ページの『gle プロパティ』](#)を参照してください。

表 171. *applygle* プロパティ

applygle プロパティ	値	プロパティの説明
enable_sql_generation	udf native	ストリーム実行中の SQL 生成オプションを設定するために使用します。データベースにプッシュバックして スポス モデラー・サーバー Scoring Adapter を使用してスコアリングするか (スコアリングアダプターがインストール済みのデータベースに接続している場合)、スポス モデラー 内でスコアリングするかを選択します。

applygmm のプロパティ

ガウス混合ノードを使用して、ガウス混合モデル・ナゲットを生成できます。このモデル・ナゲットのスクリプト名は、*applygmm* です。以下の表に示すプロパティは、バージョン 18.2.1.1 以降で使用可能です。モデル作成ノード自体をスクリプト化する方法については、[427 ページの『gmm のプロパティ』](#)を参照してください。

表 172. *applygmm* のプロパティ

applygmm のプロパティ	データ・タイプ	プロパティの説明
centers		
item_count		

表 172. *applygmm* のプロパティ (続き)

applygmm のプロパティ	データ・タイプ	プロパティの説明
total		
dimension		
components		
partition		

applykmeansnode プロパティ

K-means モデル作成ノードを使用して、K-means モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applykmeansnode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体のスクリプトの詳細は、[270 ページの『kmeansnode プロパティ』](#)を参照してください。

applyknnnode プロパティ

KNN モデル作成ノードを使用して、KNN モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyknnnode* です。モデル作成ノード自体のスクリプトの詳細は、[272 ページの『knnnode プロパティ』](#)を参照してください。

表 173. *applyknnnode* プロパティ

applyknnnode プロパティ	値	プロパティの説明
all_probabilities	フラグ	
save_distances	フラグ	

applykohonennode プロパティ

Kohonen モデル作成ノードを使用して、Kohonen モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applykohonennode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体のスクリプトの詳細は、[232 ページの『c50node プロパティ』](#)を参照してください。

applylinearnode プロパティ

線型モデル作成ノードを使用して、線型モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applylinearnode* です。モデル作成ノード自体のスクリプトの詳細は、[276 ページの『linearnode プロパティ』](#)を参照してください。

表 174. *applylinearnode* プロパティ

linear プロパティ	値	プロパティの説明
use_custom_name	フラグ	
custom_name	string	

表 174. *applylinearnode* プロパティ (続き)

linear プロパティ	値	プロパティの説明
enable_sql_generation	udf native puresql	ストリーム実行中の SQL 生成オプションを設定するために使用します。データベースにプッシュバックして SPSS® Modeler Server Scoring Adapter でスコアリングするか (スコアリングアダプターがインストール済みのデータベースに接続している場合)、SPSS Modeler 内でスコアリングするか、またはデータベースにプッシュバックして SQL でスコアリングするかを選択できます。 デフォルト値は udf です。

applylinearasnode プロパティ

Linear-AS モデル作成ノードを使用して、Linear-AS モデル ナゲットを生成できます。このモデル ナゲットのスクリプト名は *applylinearasnode* です。モデル作成ノード自体のスクリプトの詳細は、[277 ページの『linearasnode プロパティ』](#)を参照してください。

表 175. *applylinearasnode* プロパティ

applylinearasnode プロパティ	値	プロパティの説明
enable_sql_generation	udf native	デフォルト値は udf です。

applylogregnode プロパティ

ロジスティック回帰モデル作成ノードを使用して、ロジスティック回帰モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applylogregnode* です。モデル作成ノード自体のスクリプトの詳細は、[279 ページの『logregnode プロパティ』](#)を参照してください。

表 176. *applylogregnode* プロパティ

applylogregnode プロパティ	値	プロパティの説明
calculate_raw_propensities	フラグ	
calculate_conf	フラグ	
enable_sql_generation	フラグ	

applysvmnode プロパティ

LSVM モデル作成ノードを使用して、LSVM モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は *applysvmnode* です。モデル作成ノード自体をスクリプト化する方法については、[284 ページの『svmnode プロパティ』](#)を参照してください。

表 177. <i>applylsvmnode</i> プロパティ		
applylsvmnode プロパティ	値	プロパティの説明
calculate_raw_propensities	フラグ	未調整傾向スコアを計算するかどうかを指定します。
enable_sql_generation	udf native	Scoring Adapter (インストールされている場合) を使用またはインプロセスでスコアリングするか、データベースの外部でスコアリングするかを指定します。

applyneuralnetnode プロパティ

ニューラル・ネットワーク・モデル作成ノードを使用して、ニューラル・ネットワーク・モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyneuralnetnode* です。モデル作成ノード自体のスクリプトの詳細は、285 ページの『*neuralnetnode* プロパティ

注意: 機能が拡張された新しいバージョンのニューラル・ネットワーク ナゲットがこのリリースで使用できます。新しいバージョンについては次の項で説明します (*applyneuralnetwork*)。以前のバージョンは現在も使用できますが、スクリプトを更新して新しいバージョンを使用することをお勧めします。旧バージョンの詳細を参照用に記載しておりますが、それに対するサポートは今後のリリースで廃止されます。

表 178. <i>applyneuralnetnode</i> プロパティ		
applyneuralnetnode プロパティ	値	プロパティの説明
calculate_conf	フラグ	SQL 生成が有効になっている場合に利用できます。このプロパティには、生成されたツリー中の確信度計算が含まれています。
enable_sql_generation	フラグ	
nn_score_method	Difference SoftMax	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applyneuralnetworknode プロパティ

ニューラル・ネットワーク・モデル作成ノードを使用して、ニューラル・ネットワーク・モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyneuralnetworknode* です。モデル作成ノード自体のスクリプトの詳細は、*neuralnetworknode* プロパティを参照してください。

表 179. <i>applyneuralnetworknode</i> プロパティ		
applyneuralnetworknode プロパティ	値	プロパティの説明
use_custom_name	フラグ	
custom_name	string	

表 179. *applyneuralnetworknode* プロパティ (続き)

applyneuralnetworknode プロパティ	値	プロパティの説明
confidence	onProbability onIncrease	
score_category_probabilities	フラグ	
max_categories	number	
score_propensity	フラグ	
enable_sql_generation	udf native puresql	ストリーム実行中の SQL 生成オプションを設定するために使用します。データベースにプッシュバックして SPSS® Modeler Server Scoring Adapter でスコアリングするか (スコアリングアダプターがインストール済みのデータベースに接続している場合)、SPSS Modeler 内でスコアリングするか、またはデータベースにプッシュバックして SQL でスコアリングするかを選択できます。 デフォルト値は udf です。

applyocsvmnode のプロパティ

One-Class SVM ノードを使用して、One-Class SVM モデル・ナゲットを作成することができます。このモデル・ナゲットのスクリプト名は、*applyocsvmnode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体をスクリプト化する方法については、[433 ページの『ocsvmnode のプロパティ』](#)を参照してください。

applyquestnode プロパティ

QUEST モデル作成ノードを使用して、QUEST モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyquestnode* です。モデル作成ノード自体のスクリプトの詳細は、[290 ページの『questnode プロパティ』](#)を参照してください。

表 180. *applyquestnode* プロパティ

applyquestnode プロパティ	値	プロパティの説明
enable_sql_generation	Never MissingValues NoMissingValues	ルールセット実行時の SQL 生成オプションの設定に使用します。
calculate_conf	フラグ	
display_rule_id	フラグ	フィールドが 1 つスコアリング出力に追加されますが、これは各レコードを割り当てるターミナル・ノードに ID を示すためのものです。

表 180. <i>applyquestnode</i> プロパティ (続き)		
applyquestnode プロパティ	値	プロパティの説明
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applyr プロパティ

R 作成ノードを使用して、R モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyr* です。モデル作成ノード自体のスクリプトの詳細は、[231 ページの『buildr プロパティ』](#)を参照してください。

表 181. <i>applyr</i> プロパティ		
applyr プロパティ	値	プロパティの説明
score_syntax	string	モデル・スコアリング用の R スクリプト・シンタックス。
convert_flags	StringsAndDoubles LogicalValues	フラグ型フィールドを変換するためのオプション。
convert_datetime	フラグ	日付形式または日付/時刻形式の変数を R の日付/時刻形式に変換するためのオプション。
convert_datetime_class	POSIXct POSIXlt	日付形式または日付/時刻形式の変数のうち、どの形式の変数を変換するかを指定するためのオプション。
convert_missing	フラグ	欠損値を R NA 値に変換するオプションです。
use_batch_size	フラグ	バッチ処理を使用可能にします
batch_size	integer	各バッチに含めるデータレコードの数を指定します

applyrandomtrees プロパティ

Random Trees モデル作成ノードを使用して、Random Trees モデル ナゲットを生成できます。このモデルナゲットのスクリプト名は *applyrandomtrees* です。モデル作成ノード自体をスクリプト化する方法については、[293 ページの『randomtrees プロパティ』](#)を参照してください。

表 182. <i>applyrandomtrees</i> プロパティ		
applyrandomtrees プロパティ	値	プロパティの説明
calculate_conf	フラグ	このプロパティには、生成されたツリー中の確信度計算が含まれています。

表 182. *applyrandomtrees* プロパティ (続き)

applyrandomtrees プロパティ	値	プロパティの説明
enable_sql_generation	udf native	ストリーム実行中の SQL 生成オプションを設定するために使用します。データベースにプッシュバックしてスポスマデラー・サーバー Scoring Adapter を使用してスコアリングするか (スコアリングアダプタがインストール済みのデータベースに接続している場合)、スポスマデラー内でスコアリングするかを選択します。

applyregressionnode プロパティ

線型モデル作成ノードを使用して、線型モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyregressionnode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体のスクリプトの詳細は、[295 ページの『regressionnode プロパティ』](#)を参照してください。

applyselflearningnode プロパティ

自己学習応答モデル (SLRM) モデル作成ノードを使用して、SLRM モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applyselflearningnode* です。モデル作成ノード自体のスクリプトの詳細は、[299 ページの『slrmnode プロパティ』](#)を参照してください。

表 183. *applyselflearningnode* プロパティ

applyselflearningnode プロパティ	値	プロパティの説明
max_predictions	number	
randomization	number	
scoring_random_seed	number	
sort	ascending descending	高いスコアまたは低いスコアのどちらを持つオファーが最初に表示されるかを指定します。
model_reliability	フラグ	「設定」タブでモデルの信頼性を考慮します。

applysequencenode プロパティ

シーケンス・モデル作成ノードを使用して、シーケンス・モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applysequencenode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体のスクリプトの詳細は、[297 ページの『sequencenode プロパティ』](#)を参照してください。

applysvmnode プロパティ

SVM モデル作成ノードを使用して、SVM モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applysvmnode* です。モデル作成ノード自体のスクリプトの詳細は、[307 ページの『svmnode プロパティ』](#)を参照してください。

表 184. <i>applysvmnode</i> プロパティ		
applysvmnode プロパティ	値	プロパティの説明
all_probabilities	フラグ	
calculate_raw_propensities	フラグ	
calculate_adjusted_propensities	フラグ	

applystpnode プロパティ

STP モデル作成ノードを使用して、関連するモデル ナゲットを生成することができます。このモデル ナゲットにより、出力ビューアにモデル出力が表示されます。このモデル ナゲットのスクリプト名は *applystpnode* です。モデル作成ノード自体をスクリプト化する方法については、[300 ページの『stpnode プロパティ』](#)を参照してください。

表 185. <i>applystpnode</i> プロパティ		
applystpnode プロパティ	データ・タイプ	プロパティの説明
uncertainty_factor	<i>Boolean</i>	最小値は 0、最大値は 100 です。

applytcmnode プロパティ

時間的因果モデリング (TCM) モデル作成ノードを使用して、TCM モデル ナゲットを生成できます。このモデル ナゲットのスクリプト名は、*applytcmnode* です。モデル作成ノード自体をスクリプト化する方法については、[308 ページの『tcmnode プロパティ』](#)を参照してください。

表 186. <i>applytcmnode</i> プロパティ		
applytcmnode プロパティ	値	プロパティの説明
ext_future	<i>Boolean</i>	
ext_future_num	<i>integer</i>	
noise_res	<i>Boolean</i>	
conf_limits	<i>Boolean</i>	
target_fields	リスト	
target_series	リスト	

applyts プロパティ

時系列モデル作成ノードを使用して、時系列モデル ナゲットを生成することができます。このモデル ナゲットのスクリプト名は、*applyts* です。モデル作成ノード自体をスクリプト化する方法については、[313 ページの『ts プロパティ』](#)を参照してください。

表 187. <i>applyts</i> プロパティ		
applyts プロパティ	値	プロパティの説明
extend_records_into_future	<i>Boolean</i>	
ext_future_num	<i>integer</i>	
compute_future_values_input	<i>Boolean</i>	
forecastperiods	<i>integer</i>	

表 187. *applyts* プロパティ (続き)

applyts プロパティ	値	プロパティの説明
noise_res	<i>Boolean</i>	
conf_limits	<i>Boolean</i>	
target_fields	リスト	
target_series	リスト	
includeTargets	フィールド	

applytimeseriesnode プロパティ (廃止)

時系列モデル作成ノードを使用して、時系列モデルナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applytimeseriesnode* です。モデル作成ノード自体のスクリプトの詳細は、[321 ページ](#)の『*timeseriesnode* プロパティ (廃止)』を参照してください。

表 188. *applytimeseriesnode* プロパティ

applytimeseriesnode プロパティ	値	プロパティの説明
calculate_conf	フラグ	
calculate_residuals	フラグ	

applytreeas プロパティ

Tree-AS モデル作成ノードを使用して、Tree-AS モデルナゲットを生成できます。このモデルナゲットのスクリプト名は *applytreeas* です。モデル作成ノード自体をスクリプト化する方法については、[324 ページ](#)の『*treeas* プロパティ』を参照してください。

表 189. *applytreeas* プロパティ

applytreeas プロパティ	値	プロパティの説明
calculate_conf	フラグ	このプロパティには、生成されたツリー中の確信度計算が含まれています。
display_rule_id	フラグ	フィールドが 1 つスコアリング出力に追加されますが、これは各レコードを割り当てるターミナル・ノードに ID を示すためのものです。
enable_sql_generation	udf native	ストリーム実行中の SQL 生成オプションを設定するために使用します。データベースにプッシュバックしてスポンモデラー・サーバー Scoring Adapter を使用してスコアリングするか (スコアリングアダプタがインストール済みのデータベースに接続している場合)、スポンモデラー内でスコアリングするかを選択します。

applytwostepnode プロパティ

TwoStep モデル作成ノードを使用して、TwoStep モデル・ナゲットを生成することができます。このモデル・ナゲットのスクリプト名は、*applytwostepnode* です。このモデル・ナゲットの他のプロパティはあ

りません。モデル作成ノード自体のスキプトの詳細は、[326 ページの『twostepnode プロパティ』](#)を参照してください。

applytwostepAS のプロパティ

TwoStep AS モデル作成ノードを使用して、TwoStep AS モデル ナゲットを生成することができます。このモデル ナゲットのスキプト名は *applytwostepAS* です。モデル作成ノード自体のスキプトの詳細は、[327 ページの『twostepAS のプロパティ』](#)を参照してください。

表 190. <i>applytwostepAS</i> のプロパティ		
<i>applytwostepAS</i> のプロパティ	値	プロパティの説明
enable_sql_generation	udf native	ストリーム実行中の SQL 生成オプションを設定するために使用します。データベースにプッシュバックして SPSS® Modeler Server Scoring Adapter でスコアリングするか (スコアリングアダプターがインストール済みのデータベースに接続している場合)、SPSS Modeler 内でスコアリングするかを選択できます。 デフォルト値は udf です。

applyxgboosttreenode のプロパティ

XGBoost ツリー・ノードを使用して、XGBoost Tree モデル・ナゲットを生成できます。このモデル・ナゲットのスキプト名は、*applyxgboosttreenode* です。以下の表に示すプロパティは 18.2.1.1 で追加されたものです。モデル作成ノード自体をスキプト化する方法については、[441 ページの『xgboosttreenode のプロパティ』](#)を参照してください。

表 191. <i>applyxgboosttreenode</i> のプロパティ		
<i>applyxgboosttreenode</i> のプロパティ	データ・タイプ	プロパティの説明
use_model_name		
model_name		

applyxgboostlinearnode のプロパティ

XGBoost Linear ノードを使用して、XGBoost Linear モデル・ナゲットを生成できます。このモデル・ナゲットのスキプト名は、*applyxgboostlinearnode* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体をスキプト化する方法については、[440 ページの『xgboostlinearnode のプロパティ』](#)を参照してください。

hdbscannugget のプロパティ

HDBSCAN ノードを使用して、HDBSCAN モデル・ナゲットを生成できます。このモデル・ナゲットのスキプト名は、*hdbscannugget* です。このモデル・ナゲットの他のプロパティはありません。モデル作成ノード自体をスキプト化する方法については、[428 ページの『hdbscannode のプロパティ』](#)を参照してください。

kdeapply のプロパティ

KDE モデル作成ノードを使用して、KDE モデル・ナゲットを生成できます。このモデル・ナゲットのスクリプト名は、kdeapply です。モデル作成ノード自体をスクリプト化する方法については、[429 ページの『kdemodel のプロパティ』](#)を参照してください。

表 192. kdeapply のプロパティ		
kdeapply のプロパティ	データ・タイプ	プロパティの説明
outLogDensity バージョン 18.2.1.1 以降は out_log_density に名前変更 されています。	Boolean	出力の対数密度の値を含めるには True を 指定し、除外するには False を指定しま す。デフォルトは False です。

第 15 章 データベース・モデル作成ノードのプロパティ

IBM スポス モデラー は、データベース・ベンダーから入手できる、Microsoft SQL Server Analysis Services、Oracle Data Mining、InfoSphere Warehouse、IBM Netezza[®]Analytics 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。IBM スポス モデラー ネイティブ・データベース・アルゴリズムを使用して、アプリケーション内からのモデルの構築およびスコアリングがすべて可能です。データベース・モデルは、このセクションで説明するプロパティを使用してスクリプトで作成および処理することも可能です。

例えば、次のスクリプトの引用は、IBM スポス モデラー スクリプト・インターフェースを使用した Microsoft ディジジョン・ツリー・モデルの作成を示します。

```
stream = modeler.script.stream()
msbuilder = stream.createAt("mstreenode", "MSBuilder", 200, 200)

msbuilder.setPropertyValue("analysis_server_name", 'localhost')
msbuilder.setPropertyValue("analysis_database_name", 'TESTDB')
msbuilder.setPropertyValue("mode", 'Expert')
msbuilder.setPropertyValue("datasource", 'LocalServer')
msbuilder.setPropertyValue("target", 'Drug')
msbuilder.setPropertyValue("inputs", ['Age', 'Sex'])
msbuilder.setPropertyValue("unique_field", 'IDX')
msbuilder.setPropertyValue("custom_fields", True)
msbuilder.setPropertyValue("model_name", 'MSDRUG')

typenode = stream.findByType("type", None)
stream.link(typenode, msbuilder)
results = []
msbuilder.run(results)
msapplier = stream.createModelApplierAt(results[0], "Drug", 200, 300)
tablenode = stream.createAt("table", "Results", 300, 300)
stream.linkBetween(msapplier, typenode, tablenode)
msapplier.setPropertyValue("sql_generate", True)
tablenode.run([])
```

Microsoft モデル作成ノードのプロパティ

Microsoft モデル作成ノードのプロパティ

共通のプロパティ

次のプロパティは、Microsoft データベース・モデル作成ノードに共通です。

表 193. 共通の Microsoft ノード・プロパティ		
共通の Microsoft ノード・プロパティ	値	プロパティの説明
analysis_database_name	string	Analysis Services データベースの名前。
analysis_server_name	string	Analysis Services ホストの名前。
use_transactional_data	フラグ	入力データがテーブル形式またはトランザクション形式かを指定します。
inputs	リスト	テーブル形式の入力フィールド。

表 193. 共通の Microsoft ノード・プロパティ (続き)		
共通の Microsoft ノード・プロパティ	値	プロパティの説明
target	フィールド	予測フィールド (MS Clustering または Sequence Clustering ノードには該当しない)。
unique_field	フィールド	キー・フィールド。
msas_parameters	構造化	アルゴリズム・パラメーター。詳しくは、トピック 355 ページの『 アルゴリズム・パラメーター 』を参照してください。
with_drillthrough	フラグ	「ドリルスルーあり」オプション。

MS ディジション・ツリー

mstreenode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Microsoft プロパティを参照してください。

MS クラスターリング

msclusternode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Microsoft プロパティを参照してください。

MS アソシエーション・ルール

次のプロパティは、msassocnode タイプのノードで使用できます。

表 194. msassocnode プロパティ		
msassocnode プロパティ	値	プロパティの説明
id_field	フィールド	データの各トランザクションを特定します。
trans_inputs	リスト	トランザクションデータの入力フィールド。
transactional_target	フィールド	予測データ (トランザクション・データ)。

MS Naive Bayes

msbayesnode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Microsoft プロパティを参照してください。

MS Linear Regression

msregressionnode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Microsoft プロパティを参照してください。

MS ニューラル・ネットワーク

msneuralnetworknode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Microsoft プロパティを参照してください。

MS Logistic Regression

mslogisticnode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Microsoft プロパティを参照してください。

MS Time Series

mstimeseriesnode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Microsoft プロパティを参照してください。

MS Sequence Clustering

次のプロパティは、mssequenceclusternode タイプのノードで使用できます。

表 195. mssequenceclusternode プロパティ		
mssequenceclusternode プロパティ	値	プロパティの説明
id_field	フィールド	データの各トランザクションを特定します。
input_fields	リスト	トランザクションデータの入力フィールド。
sequence_field	フィールド	シーケンス ID。
target_field	フィールド	予測フィールド (テーブル形式データ)。

アルゴリズム・パラメーター

各 Microsoft データベース・モデル・タイプには、msas_parameters プロパティを使用して設定できる特定のパラメーターがあります。以下に例を示します。

```
stream = modeler.script.stream()
msregressionnode = stream.findByType("msregression", None)
msregressionnode.setPropertyValue("msas_parameters",
[["MAXIMUM_INPUT_ATTRIBUTES", 255],
["MAXIMUM_OUTPUT_ATTRIBUTES", 255]])
```

これらのパラメーターは SQL Server から取得されます。各ノードに関連するパラメーターを見るには

1. キャンバスにデータベース入力ノードを配置します。
2. データベース入力ノードを開きます。
3. 「データソース」 ドロップダウン・リストから有効なソースを選択します。
4. 「テーブル名」 リストから有効なテーブルを選択します。
5. 「OK」 をクリックして、データベース入力ノードを閉じます。
6. プロパティを一覧表示したい Microsoft データベース・モデル作成ノードを追加します。
7. データベース・モデル作成ノードを開きます。
8. 「エクスポート」 タブを選択します。

このノードの使用できる msas_parameters プロパティが表示されます。

Microsoft モデル・ナゲットのプロパティ

Microsoft データベース・モデル作成ノードを使用して作成されるモデル・ナゲットのプロパティを、次に示します。

MS ディジジョン・ツリー

表 196. MS ディジジョン・ツリーのプロパティ

appliednode プロパティ	値	説明
analysis_database_name	string	このノードは、ストリームの中で直接スコアされます。 このプロパティは Analysis Services データベース名の識別に使用します。
analysis_server_name	string	Analysis サーバー・ホストの名前
datasource	string	SQL Server の ODBC データ・ソース 名 (DSN) の名前
sql_generate	フラグ udf	SQL 生成を有効にします。

MS Linear Regression

表 197. MS Linear Regression のプロパティ

appliednode プロパティ	値	説明
analysis_database_name	string	このノードは、ストリームの中で直接スコアされます。 このプロパティは Analysis Services データベース名の識別に使用します。
analysis_server_name	string	Analysis サーバー・ホストの名前

MS ニューラル・ネットワーク

表 198. MS Neural Network のプロパティ

appliednode プロパティ	値	説明
analysis_database_name	string	このノードは、ストリームの中で直接スコアされます。 このプロパティは Analysis Services データベース名の識別に使用します。
analysis_server_name	string	Analysis サーバー・ホストの名前

MS Logistic Regression

表 199. MS Logistic Regression のプロパティ		
applieslogisticnode プロパティ	値	説明
analysis_database_name	string	このノードは、ストリームの中で直接スコアされます。 このプロパティは Analysis Services データベース名の識別に使用します。
analysis_server_name	string	Analysis サーバー・ホストの名前

MS Time Series

表 200. MS Time Series のプロパティ		
appliestimeseriesnode プロパティ	値	説明
analysis_database_name	string	このノードは、ストリームの中で直接スコアされます。 このプロパティは Analysis Services データベース名の識別に使用します。
analysis_server_name	string	Analysis サーバー・ホストの名前
start_from	new_prediction historical_prediction	将来の予測を行うか過去の予測を行うかを指定します。
new_step	number	将来の予測の開始時間を定義します。
historical_step	number	過去の予測の開始時間を定義します。
end_step	number	予測の終了時間を定義します。

MS Sequence Clustering

表 201. MS Sequence Clustering のプロパティ		
appliessequenceclusternode プロパティ	値	説明
analysis_database_name	string	このノードは、ストリームの中で直接スコアされます。 このプロパティは Analysis Services データベース名の識別に使用します。
analysis_server_name	string	Analysis サーバー・ホストの名前

Oracle モデル作成ノードのプロパティ

Oracle モデル作成ノードのプロパティ

次のプロパティは、各 Oracle データベース・モデル作成ノードに共通です。

表 202. Oracle ノードの共通プロパティ		
一般的な Oracle ノードのプロパティ	値	プロパティの説明
target	フィールド	
inputs	フィールドのリスト	
partition	フィールド	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。
datasource		
username		
password		
epassword		
use_model_name	フラグ	
model_name	string	ユーザーが指定する新規モデル名。
use_partitioned_data	フラグ	区分フィールドが定義される場合、このオプションは学習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。
unique_field	フィールド	
auto_data_prep	フラグ	Oracle データの自動準備機能を有効化または無効化します (11g データベースのみ)。
costs	構造化	構造化プロパティ、使用形式： [[drugA drugB 1.5] [drugA drugC 2.1]]。[] 内の引数は実際の予測コストです。
mode	Simple Expert	Simple に設定されている場合、個々のノード・プロパティに記述されているように、特定のプロパティは無視されます。
use_prediction_probability	フラグ	
prediction_probability	string	
use_prediction_set	フラグ	

Oracle Naive Bayes

次のプロパティは、oranbnode タイプのノードで使用できます。

表 203. oranbnode プロパティ		
oranbnode プロパティ	値	プロパティの説明
singleton_threshold	number	0.0–1.0.*
pairwise_threshold	number	0.0–1.0.*

表 203. oranbnode プロパティ (続き)

oranbnode プロパティ	値	プロパティの説明
priors	Data Equal Custom	
custom_priors	構造化	構造化プロパティ、使用形式: set :oranbnode.custom_priors = [[drugA 1][drugB 2][drugC 3][drugX 4][drugY 5]]

* mode が Simple に設定されている場合、プロパティは無視されます。

Oracle Adaptive Bayes

次のプロパティは、oraabnnode タイプのノードで使用できます。

表 204. oraabnnode プロパティ

oraabnnode プロパティ	値	プロパティの説明
model_type	SingleFeature MultiFeature NaiveBayes	
use_execution_time_limit	フラグ	*
execution_time_limit	integer	値は 1 以上でなければなりません。*
max_naive_bayes_predictors	integer	値は 1 以上でなければなりません。*
max_predictors	integer	値は 1 以上でなければなりません。*
priors	Data Equal Custom	
custom_priors	構造化	構造化プロパティ、使用形式: set :oraabnnode.custom_priors = [[drugA 1][drugB 2][drugC 3][drugX 4][drugY 5]]

* mode が Simple に設定されている場合、プロパティは無視されます。

Oracle Support Vector Machines

次のプロパティは、orasvmnode タイプのノードで使用できます。

表 205. orasvmnode プロパティ

orasvmnode プロパティ	値	プロパティの説明
active_learning	Enable Disable	
kernel_function	Linear Gaussian System	
normalization_method	zscore minmax none	
kernel_cache_size	<i>integer</i>	Gaussian カーネル専用。値は 1 以上でなければなりません。*
convergence_tolerance	<i>number</i>	値は 1 以上でなければなりません。*
use_standard_deviation	フラグ	Gaussian カーネル専用。*
standard_deviation	<i>number</i>	値は 1 以上でなければなりません。*
use_epsilon	フラグ	回帰モデルのみです。*
epsilon	<i>number</i>	値は 1 以上でなければなりません。*
use_complexity_factor	フラグ	*
complexity_factor	<i>number</i>	*
use_outlier_rate	フラグ	単一バリエーションのみです。*
outlier_rate	<i>number</i>	単一バリエーションのみです。0.0–1.0.*
weights	Data Equal Custom	
custom_weights	構造化	構造化プロパティ、使用形式： set :orasvmnode.custom_weights = [[drugA 1][drugB 2][drugC 3][drugX 4][drugY 5]]

* mode が Simple に設定されている場合、プロパティは無視されます。

Oracle 一般化線型モデル

次のプロパティは、`oraglmnode` タイプのノードで使用できます。

表 206. <i>oraglmnode</i> プロパティ		
oraglmnode プロパティ	値	プロパティの説明
normalization_method	zscore minmax none	
missing_value_handling	ReplaceWithMean UseCompleteRecords	
use_row_weights	フラグ	*
row_weights_field	フィールド	*
save_row_diagnostics	フラグ	*
row_diagnostics_table	<i>string</i>	*
coefficient_confidence	<i>number</i>	*
use_reference_category	フラグ	*
reference_category	<i>string</i>	*
ridge_regression	Auto Off On	*
parameter_value	<i>number</i>	*
vif_for_ridge	フラグ	*

* `mode` が `Simple` に設定されている場合、プロパティは無視されます。

Oracle ディジジョン・ツリー

次のプロパティは、`oradecisiontreenode` タイプのノードで使用できます。

表 207. <i>oradecisiontreenode</i> プロパティ		
oradecisiontreenode プロパティ	値	プロパティの説明
use_costs	フラグ	
impurity_metric	Entropy Gini	
term_max_depth	<i>integer</i>	2-20.*
term_minpct_node	<i>number</i>	0.0-10.0.*

表 207. oradecisiontreenode プロパティ (続き)

oradecisiontreenode プロパティ	値	プロパティの説明
term_minpct_split	number	0.0–20.0.*
term_minrec_node	integer	値は 1 以上でなければなりません。*
term_minrec_split	integer	値は 1 以上でなければなりません。*
display_rule_ids	フラグ	*

*mode が Simple に設定されている場合、プロパティは無視されます。

Oracle O-Cluster

次のプロパティは、oraoclusternode タイプのノードで使用できます。

表 208. oraoclusternode プロパティ

oraoclusternode プロパティ	値	プロパティの説明
max_num_clusters	integer	値は 1 以上でなければなりません。*
max_buffer	integer	値は 1 以上でなければなりません。*
sensitivity	number	0.0–1.0.*

*mode が Simple に設定されている場合、プロパティは無視されます。

Oracle KMeans

次のプロパティは、orakmeansnode タイプのノードで使用できます。

表 209. orakmeansnode プロパティ

orakmeansnode プロパティ	値	プロパティの説明
num_clusters	integer	値は 1 以上でなければなりません。*
normalization_method	zscore minmax none	
distance_function	Euclidean Cosine	
iterations	integer	0–20.*
conv_tolerance	number	0.0–0.5.*
split_criterion	Variance Size	デフォルトは Variance です。*

表 209. orakmeansnode プロパティ (続き)		
orakmeansnode プロパティ	値	プロパティの説明
num_bins	integer	値は 1 以上でなければなりません。*
block_growth	integer	1-5.*
min_pct_attr_support	number	0.0-1.0.*

* mode が Simple に設定されている場合、プロパティは無視されます。

Oracle NMF

次のプロパティは、oranmfnode タイプのノードで使用できます。

表 210. oranmfnode プロパティ		
oranmfnode プロパティ	値	プロパティの説明
normalization_method	minmax	
	none	
use_num_features	フラグ	*
num_features	integer	0-1. デフォルト値はアルゴリズムによってデータから推定されます。
random_seed	number	*
num_iterations	integer	0-500.*
conv_tolerance	number	0.0-0.5.*
display_all_features	フラグ	*

* mode が Simple に設定されている場合、プロパティは無視されます。

Oracle Apriori

次のプロパティは、oraapriorinode タイプのノードで使用できます。

表 211. oraapriorinode プロパティ		
oraapriorinode プロパティ	値	プロパティの説明
content_field	フィールド	
id_field	フィールド	
max_rule_length	integer	2-20.
min_confidence	number	0.0-1.0.
min_support	number	0.0-1.0.
use_transactional_data	フラグ	

Oracle 最小記述長 (MDL)

oramdlnode タイプのノードには、特定のプロパティが定義されていません。このセクションの冒頭にある共通 Oracle プロパティを参照してください。

Oracle Attribute Importance (AI)

次のプロパティは、oraainode タイプのノードで使用できます。

表 212. oraainode プロパティ		
oraainode プロパティ	値	プロパティの説明
custom_fields	フラグ	真 (true) の場合は、現在のノードのターゲット、入力、その他フィールドなどを指定することができます。偽 (false) の場合は、上流のデータ型ノードから現在の設定が使用されます。
selection_mode	ImportanceLevel ImportanceValue TopN	
select_important	フラグ	selection_mode が ImportanceLevel に設定されているときに、重要なフィールドを選択するかどうかを指定します。
important_label	string	「重要」ランクのラベルを指定します。
select_marginal	フラグ	selection_mode が ImportanceLevel に設定されているときに、境界フィールドを選択するかどうかを指定します。
marginal_label	string	「境界」ランクのラベルを指定します。
important_above	number	0.0–1.0.
select_unimportant	フラグ	selection_mode が ImportanceLevel に設定されているときに、重要でないフィールドを選択するかどうかを指定します。
unimportant_label	string	「非重要」ランクのラベルを指定します。
unimportant_below	number	0.0–1.0.
importance_value	number	selection_mode が ImportanceValue に設定されているときに、使用する分割値を指定します。0 から 100 の値を指定します。
top_n	number	selection_mode が TopN に設定されているときに、使用する分割値を指定します。0 から 1000 の値を指定します。

Oracle モデル・ナゲットのプロパティ

Oracle ノードを使用して作成されるモデル・ナゲットのプロパティを、次に示します。

Oracle Naive Bayes

applyoranbnode タイプのノードには、特定のプロパティが定義されていません。

Oracle Adaptive Bayes

applyoraabnnode タイプのノードには、特定のプロパティが定義されていません。

Oracle Support Vector Machines

applyorasvmnode タイプのノードには、特定のプロパティが定義されていません。

Oracle ディシジョン・ツリー

次のプロパティは、applyoradecisiontreenode タイプのノードで使用できます。

表 213. applyoradecisiontreenode プロパティ		
applyoradecisiontreenode プロパティ	値	プロパティの説明
use_costs	フラグ	
display_rule_ids	フラグ	

Oracle O-Cluster

applyoraoclusternode タイプのノードには、特定のプロパティが定義されていません。

Oracle KMeans

applyorakmeansnode タイプのノードには、特定のプロパティが定義されていません。

Oracle NMF

次のプロパティは、applyoranmfnode タイプのノードで使用できます。

表 214. applyoranmfnode プロパティ		
applyoranmfnode プロパティ	値	プロパティの説明
display_all_features	フラグ	

Oracle Apriori

このモデル・ナゲットはスクリプトに適用できません。

Oracle MDL

このモデル・ナゲットはスクリプトに適用できません。

IBM NetezzaAnalytics モデル作成ノードのプロパティ

Netezza モデル作成ノードのプロパティ

次のプロパティは、各 IBM Netezza データベース・モデル作成ノードに共通です。

表 215. 共通の Netezza ノード・プロパティ		
共通の Netezza ノード・プロパティ	値	プロパティの説明
custom_fields	フラグ	真 (true) の場合は、現在のノードのターゲット、入力、その他フィールドなどを指定することができます。偽 (false) の場合は、上流のデータ型ノードから現在の設定が使用されます。

表 215. 共通の Netezza ノード・プロパティ (続き)

共通の Netezza ノード・プロパティ	値	プロパティの説明
inputs	[field1 ... fieldN]	モデルで使用される入力または予測変数フィールド。
target	フィールド	対象フィールド (連続型またはカテゴリー型)。
record_id	フィールド	一意のレコード ID として使用されるフィールド。
use_upstream_connection	フラグ	true (デフォルト) の場合、上流のノードで指定された接続の詳細。move_data_to_connection が指定されている場合は使用されません。
move_data_connection	フラグ	true の場合、データは connection に指定されたデータベースに移動します。use_upstream_connection が指定されている場合は使用されません。
connection	構造化	モデルが保存される Netezza データベースの接続文字列。構造化プロパティ、使用形式： <pre>['odbc' '<dsn>' '<username>' '<psw>' '<catname>' '<conn_attribs>' [true false]]</pre> ここで、 <dsn> はデータ・ソース名です。 <username> と <psw> は、データベースのユーザー名とパスワードです。 <catname> はカタログ名です。 <conn_attribs> は接続の属性です。 true false は、パスワードが必要かどうかを示します。
table_name	string	モデルが保存されるデータベース・テーブルの名前。
use_model_name	フラグ	true の場合、model_name によって指定された名前をモデルの名前として使用します。そうでない場合、モデル名はシステムによって作成されます。
model_name	string	ユーザーが指定する新規モデル名。
include_input_fields	フラグ	true の場合、すべての入力フィールドを下流に渡します。そうでない場合、record_id とモデルによって生成されたフィールドのみが渡されます。

Netezza ディシジョン・ツリー

次のプロパティは、netezzadectreenode タイプのノードで使用できます。

表 216. netezzadectreenode プロパティ

netezzadectreenode プロパティ	値	プロパティの説明
impurity_measure	Entropy Gini	ツリーの分割に最も良い場所を評価するのに使用される、不純度の測定。
max_tree_depth	integer	ツリーが成長可能な最大レベル数。デフォルトは 62 です (可能な最大値)。
min_improvement_splits	number	分割が発生する不純度の改善の最小値。デフォルトは 0.01 です。
min_instances_split	integer	分割が発生する前に残る分割されていないレコードの最小数。デフォルトは 2 です (可能な最小値)。
weights	構造化	クラスの相対的重み。構造化プロパティ、使用形式: <pre>set :netezza_dectree.weights = [[drugA 0.3][drugB 0.6]]</pre> デフォルトの重みはすべてのクラスで 1 です。
pruning_measure	Acc wAcc	デフォルトは Acc (精度) です。wAcc (重み付き精度) は、剪定を適用する際にクラスの重みを考慮します。
prune_tree_options	allTrainingData partitionTrainingData useOtherTable	デフォルトでは、allTrainingData を使用してモデルの精度を推定します。partitionTrainingData を使用して、使用する学習データの割合を、useOtherTable を使用して指定したデータベース・テーブルの学習データ・セットを使用します。
perc_training_data	number	prune_tree_options が partitionTrainingData に設定されている場合、学習に使用するデータの割合を指定します。
prune_seed	integer	prune_tree_options が partitionTrainingData に設定されている場合、分析結果を再現に使用するランダム・シード。デフォルトは 1 です。
pruning_table	string	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。

表 216. netezzadectreenode プロパティ (続き)		
netezzadectreenode プロパティ	値	プロパティの説明
compute_probabilities	フラグ	true の場合、予測フィールドのほか、確信度 (確率) フィールドを生成します。

Netezza K-Means

次のプロパティは、netezzakmeansnode タイプのノードで使用できます。

表 217. netezzakmeansnode プロパティ		
netezzakmeansnode プロパティ	値	プロパティの説明
distance_measure	Euclidean Manhattan Canberra maximum	データ・ポイント間の距離を測定する方法。
num_clusters	integer	作成するクラスター数。デフォルトは 3。
max_iterations	integer	モデルの学習を停止する前のアルゴリズムの反復数。デフォルトは 5。
rand_seed	integer	分析結果の反復に使用するランダム・シード。デフォルトは 12345。

Netezza ベイズ・ネットワーク

次のプロパティは、netezزابayesnode タイプのノードで使用できます。

表 218. netezزابayesnode プロパティ		
netezزابayesnode プロパティ	値	プロパティの説明
base_index	integer	内部管理の最初の入力フィールドに割り当てられる数値の識別子。デフォルトは 777。
sample_size	integer	属性の値が非常に大きい場合に使用するサンプルのサイズ。デフォルトは 10,000。
display_additional_information	フラグ	true の場合、メッセージのダイアログ・ボックスに追加の進捗状況の情報を表示します。
type_of_prediction	best neighbors nn-neighbors	使用する予測アルゴリズムの種類: 最適 (相関度が最も高い近傍)、近傍 (近傍の重み付き予測)、NN 近傍 (null 以外の近傍)。

Netezza Naive Bayes

次のプロパティは、netezzanaivebayesnode タイプのノードで使用できます。

表 219. netezzanaivebayesnode プロパティ		
netezzanaivebayesnode プロパティ	値	プロパティの説明
compute_probabilities	フラグ	true の場合、予測フィールドのほか、確信度 (確率) フィールドを生成します。
use_m_estimation	フラグ	true の場合、推定時に 0 の確立を回避する m 推定方法を使用します。

Netezza KNN

次のプロパティは、netezzaknnnode タイプのノードで使用できます。

表 220. netezzaknnnode プロパティ		
netezzaknnnode プロパティ	値	プロパティの説明
weights	構造化	重みを各クラスに割り当てる構造化プロパティ。 例: set :netezzaknnnode.weights = [[drugA 0.3][drugB 0.6]]
distance_measure	Euclidean Manhattan Canberra Maximum	データ・ポイント間の距離を測定する方法。
num_nearest_neighbors	integer	特定のケースの最近傍数。デフォルトは 3。
standardize_measurements	フラグ	true の場合、距離の値を計算する前に連続型入力フィールドの測定を標準化します。
use_coresets	フラグ	true の場合、大規模なデータセットに対して計算を高速化するコアセット・サンプリングを使用しています

Netezza 分裂クラスタリング

次のプロパティは、netezzaclusternode タイプのノードで使用できます。

表 221. netezzadivclusternode プロパティ		
netezzadivclusternode プロパティ	値	プロパティの説明
distance_measure	Euclidean Manhattan Canberra Maximum	データ・ポイント間の距離を測定する方法。
max_iterations	integer	モデルの学習が停止する前に、実行するアルゴリズム反復の最大回数。デフォルトは 5 です。
max_tree_depth	integer	データセットを分割することができるレベルの最大数。デフォルトは 3 です。
rand_seed	integer	分析を複製するために使用されるランダムシード。デフォルトは 12345。
min_instances_split	integer	分割可能な最小レコード数。デフォルトは 5。
level	integer	レコードをスコアリングする階層レベル。デフォルトは -1。

Netezza PCA

次のプロパティは、netezzapcanode タイプのノードで使用できます。

表 222. netezzapcanode プロパティ		
netezzapcanode プロパティ	値	プロパティの説明
center_data	フラグ	true (デフォルト) の場合、このオプションをチェックした場合、分析前にデータのセンタリングを(または「平均値減算」)を実行します。
perform_data_scaling	フラグ	true の場合、分析前にデータのスケールリングを行います。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。
force_eigensolve	フラグ	true の場合、主成分分析を計算する精度が低くなくてもより高速な方法を使用します。
pc_number	integer	データセットを減少する主要成分の数。デフォルトは 1。

Netezza 回帰ツリー

次のプロパティは、netezzaregtreenode タイプのノードで使用できます。

表 223. netezzaregtreenode プロパティ		
netezzaregtreenode プロパティ	値	プロパティの説明
max_tree_depth	integer	ルート・ノードの前にツリーが成長できるレベルの最大数。デフォルトは 10 です。

表 223. netezzaregtreenode プロパティ (続き)

netezzaregtreenode プロパティ	値	プロパティの説明
split_evaluation_measure	Variance	ツリーを分割するのに最適な場所を評価するために使用される、クラスの不純度の測定。デフォルト (現在唯一のオプション) は Variance。
min_improvement_splits	number	ツリー内に新しい分割が作成される前に純度を減少させる最小数。
min_instances_split	integer	分割可能な最小レコード数。
pruning_measure	mse r2 pearson spearman	剪定に使用する方法
prune_tree_options	allTrainingData partitionTrainingData useOtherTable	デフォルトでは、allTrainingData を使用してモデルの精度を推定します。partitionTrainingData を使用して、使用する学習データの割合を、useOtherTable を使用して指定したデータベース・テーブルの学習データ・セットを使用します。
perc_training_data	number	prune_tree_options が PercTrainingData に設定されている場合、学習に使用するデータの割合を指定します。
prune_seed	integer	prune_tree_options が PercTrainingData に設定されている場合、分析結果を再現に使用するランダム・シード。デフォルトは1です。
pruning_table	string	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。
compute_probabilities	フラグ	true の場合、割り当てられたクラスの分散が出力に含まれるべきかどうかを指定します。

Netezza 線型

次のプロパティは、netezزالineregressionnode タイプのノードで使用できます。

表 224. netezزالineregressionnode プロパティ

netezزالineregressionnode プロパティ	値	プロパティの説明
use_svd	フラグ	true の場合、元のマトリックスの代わりに特異値分解マトリックスを使用して速度と数値の精度を向上させます。
include_intercept	フラグ	true (デフォルト) の場合、ソリューションの全体の精度が向上します。
calculate_model_diagnostics	フラグ	true の場合、モデルの診断を計算します。

Netezza 時系列

次のプロパティは、netezzatimeseriesnode タイプのノードで使用できます。

表 225. netezzatimeseriesnode プロパティ

netezzatimeseriesnode プロパティ	値	プロパティの説明
time_points	フィールド	時系列の日付または時刻の値を含む入力フィールド。
time_series_ids	フィールド	時系列 ID を含むフィールド。入力に複数の時系列が含まれる場合に使用します。
model_table	フィールド	Netezza 時系列モデルが保存されるデータベース・テーブルの名前。
description_table	フィールド	時系列名および説明を含む入力テーブルの名前。
seasonal_adjustment_table	フィールド	指数平滑化または季節的傾向分解アルゴリズムによって計算された季節性調整値を保存する出力テーブル名。
algorithm_name	SpectralAnalysis または spectral ExponentialSmoothing または esmoothing ARIMA SeasonalTrendDecomposition または std	時系列モデリングに使用するアルゴリズム

表 225. *netezzatimeseriesnode* プロパティ (続き)

netezzatimeseriesnode プロパティ	値	プロパティの説明
trend_name	N A DA M DM	指数平滑化の傾向タイプ。 N - none A - 付加 DA - 付加減衰 M - 倍数 DM - 倍数減衰
seasonality_type	N A M	指数平滑化の季節性タイプ。 N - none A - 付加 M - 倍数
interpolation_method	linear cubicspline exponentialspline	使用する補間方法。
timerange_setting	SD SP	使用する時間範囲の設定。 SD - システム決定 (時系列データの全範囲を使用) SP - earliest_time および latest_time を使用したユーザー指定

表 225. *netezzatimeseriesnode* プロパティ (続き)

netezzatimeseriesnode プロパティ	値	プロパティの説明
earliest_time	<i>integer</i>	timerange_setting が SP の場合の開始値および終了値。
latest_time	日付 時間 タイム・スタンプ	形式は、time_points 値に従う必要があります。 例えば、time_points フィールドに日付が含まれる場合は、これも日付とする必要があります。 例: <pre>set NZ_DT1.timerange_setting = 'SP' set NZ_DT1.earliest_time = '1921-01-01' set NZ_DT1.latest_time = '2121-01-01'</pre>

表 225. *netezzatimeseriesnode* プロパティ (続き)

netezzatimeseriesnode プロパティ	値	プロパティの説明
arima_setting	SD SP	ARIMA アルゴリズムの設定 (algorithm_name が ARIMA に設定されている場合にのみ使用されます)。 SD - system-determined SP - user-specified arima_setting = SP の場合、以下のパラメーターを使用して季節性の値と非季節性の値を設定します。例 (非季節性のみ): <pre>set NZ_DT1.algorithm_name = 'arima'</pre> <pre>set NZ_DT1.arima_setting = 'SP'</pre> <pre>set NZ_DT1.p_symbol = 'lesseq'</pre> <pre>set NZ_DT1.p = '4'</pre> <pre>set NZ_DT1.d_symbol = 'lesseq'</pre> <pre>set NZ_DT1.d = '2'</pre> <pre>set NZ_DT1.q_symbol = 'lesseq'</pre> <pre>set NZ_DT1.q = '4'</pre>
p_symbol	less eq lesseq	ARIMA - p、d、q、sp、sd および sq パラメーターの演算子です。 less - より小さい eq - 等しい lesseq - 次の値以下
d_symbol		
q_symbol		
sp_symbol		
sd_symbol		
sq_symbol		
p	integer	ARIMA - 自己相関の非季節性の度合い。
q	integer	ARIMA - 自己相関の非季節性導出値。

表 225. *netezzatimeseriesnode* プロパティ (続き)

netezzatimeseriesnode プロパティ	値	プロパティの説明
d	<i>integer</i>	ARIMA - モデル内の移動平均の非季節性数値。
sp	<i>integer</i>	ARIMA - 自己相関の季節性の度合い。
sq	<i>integer</i>	ARIMA - 自己相関の季節性導出値。
sd	<i>integer</i>	ARIMA - モデル内の移動平均の季節性数値。
advanced_setting	SD SP	詳細設定の処理方法を決定します。 SD - system-determined SP - period、units_period および forecast_setting を使用したユーザー指定。 例: <pre>set NZ_DT1.advanced_setting = 'SP' set NZ_DT1.period = 5 set NZ_DT1.units_period = 'd'</pre>
period	<i>integer</i>	units_period と組み合わせて指定した季節性サイクルの長さ。スペクトル解析には適用できません。

表 225. *netezzatimeseriesnode* プロパティ (続き)

netezzatimeseriesnode プロパティ	値	プロパティの説明
<code>units_period</code>	<code>ms</code> <code>s</code> <code>min</code> <code>h</code> <code>d</code> <code>wk</code> <code>q</code> <code>y</code>	<code>period</code> が表現される単位。 <code>ms</code> - ミリ秒 <code>s</code> - 秒 <code>min</code> - 分 <code>h</code> - 時 <code>d</code> - 日 <code>wk</code> - 週 <code>q</code> - quarters <code>y</code> - years 例えば、1 週間は <code>period</code> に 1、<code>units_period</code> に <code>wk</code> を指定します。
<code>forecast_setting</code>	<code>forecasthorizon</code> <code>forecasttimes</code>	予測の作成方法を指定します。
<code>forecast_horizon</code>	<code>integer</code> 日付 時間 タイム・スタンプ	<code>forecast_setting</code> = <code>forecasthorizon</code> である場合、予測の終点の値を指定します。 形式は、<code>time_points</code> 値に従う必要があります。 例えば、<code>time_points</code> フィールドに日付が含まれる場合は、これも日付とする必要があります。
<code>forecast_times</code>	<code>integer</code> 日付 時間 タイム・スタンプ	<code>forecast_setting</code> = <code>forecasttimes</code> の場合、予測を作成するために使用する値を指定します。 形式は、<code>time_points</code> 値に従う必要があります。 例えば、<code>time_points</code> フィールドに日付が含まれる場合は、これも日付とする必要があります。

表 225. netezzatimeseriesnode プロパティ (続き)

netezzatimeseriesnode プロパティ	値	プロパティの説明
include_history	フラグ	過去の値を出力に含めるかどうかを示します。
include_interpolated_values	フラグ	補間された値を出力に含めるかどうかを示します。 include_history が false の場合は使用されません。

Netezza 一般化線型

次のプロパティは、netezzaglmnode タイプのノードで使用できます。

表 226. netezzaglmnode プロパティ

netezzaglmnode プロパティ	値	プロパティの説明
dist_family	bernoulli gaussian poisson negativebinomial wald gamma	分布のタイプ。デフォルトは bernoulli です。
dist_params	number	使用する分布パラメーター値。 distribution が Negativebinomial の場合のみ適用されます。
trials	integer	distribution が Binomial の場合のみ適用されます。ターゲット応答が一連の試行が発生するさまざまなイベントの場合、target フィールドにはイベント数、trials フィールドには試行回数が含まれます。
model_table	フィールド	Netezza 一般化線型モデルが保存されるデータベース・テーブルの名前。
maxit	integer	アルゴリズムが実行できる反復の最大回数。デフォルトは 20 です。
eps	number	アルゴリズムが適合度モデルの検索を停止する最大誤差の値 (科学的表記)。デフォルトは -3、つまり 1E-3 または 0.001 です。

表 226. netezzaglmnode プロパティ (続き)

netezzaglmnode プロパティ	値	プロパティの説明
tol	<i>number</i>	誤差が 0 として扱われる値 (科学的表記)。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。
link_func	identity inverse invnegative invsquare sqrt power oddspower log clog loglog cloglog logit probit gaussit cauchit canbinom cangeom cannegbinom	使用するリンク関数。デフォルトは logit です。
link_params	<i>number</i>	使用するリンク関数パラメーター値。link_function が power または oddspower の場合のみ適用されます。

表 226. netezzaglmnode プロパティ (続き)		
netezzaglmnode プロパティ	値	プロパティの説明
interaction	[[[colnames1],[levels1]], [[colnames2],[levels2]], ...,[列名 N],[レベル N]],]	フィールド間の交互作用を指定します。colnames は、入力フィールドのリストです。また、各フィールドの level は常に 0 です。 例: [[["K","BP","Sex","K"], [0,0,0,0]], [["Age","Na"],[0,0]]]
intercept	フラグ	true の場合、モデルに定数項を含みます。

Netezza モデル・ナゲットのプロパティ

次のプロパティは、Netezza データベース・モデリング ナゲットに共通です。

表 227. Netezza モデル・ナゲットの共通プロパティ		
Netezza モデル・ナゲットの共通プロパティ	値	プロパティの説明
connection	string	モデルが保存される Netezza データベースの接続文字列。
table_name	string	モデルが保存されるデータベース・テーブルの名前。

他のモデル・ナゲットのプロパティは、対応するモデル作成ノードの場合と同じです。

モデル・ナゲットのスクリプト名は以下の通りです。

表 228. Netezza モデル・ナゲットのスクリプト名	
モデル・ナゲット	スクリプト名
ディシジョン・ツリー	applynetezzaedctreenode
K-Means	applynetezzakmeansnode
ベイズ・ネット	applynetezزابayesnode
Naive Bayes	applynetezzanaivebayesnode
KNN	applynetezzaknnnode
分裂クラスタリング	applynetezzaclusternode
PCA	applynetezzapcanode
回帰ツリー	applynetezzaregtreenode
線型	applynetezzaalineregressionnode
時系列	applynetezzatimeseriesnode
一般化線型	applynetezzaglmnode

第 16 章 出力ノードのプロパティ

出力ノードのプロパティは、ほかの種類のノードのプロパティと少し異なっています。出力ノードのプロパティは、特定のノード・オプションを参照するというよりは、参照を出力オブジェクトに格納します。このことはテーブルから値を取得して、それをストリーム・パラメーターとして設定するような場合などに役立ちます。

このセクションで、出力ノードで使えるスクリプト用のプロパティを説明します。

analysisnode プロパティ



精度分析ノードで、予測モデルの能力を評価して正確な予測を生成します。精度分析ノードでは、1つ以上のモデル・ナゲットについて、予測値と実際値をさまざまな方法で比較します。また、予測モデルを相互に比較することもできます。

例

```
node = stream.create("analysis", "My node")
# "Analysis" tab
node.setPropertyValue("coincidence", True)
node.setPropertyValue("performance", True)
node.setPropertyValue("confidence", True)
node.setPropertyValue("threshold", 75)
node.setPropertyValue("improve_accuracy", 3)
node.setPropertyValue("inc_user_measure", True)
# "Define User Measure..."
node.setPropertyValue("user_if", "@TARGET = @PREDICTED")
node.setPropertyValue("user_then", "101")
node.setPropertyValue("user_else", "1")
node.setPropertyValue("user_compute", ["Mean", "Sum"])
node.setPropertyValue("by_fields", ["Drug"])
# "Output" tab
node.setPropertyValue("output_format", "HTML")
node.setPropertyValue("full_filename", "C:/output/analysis_out.html")
```

表 229. analysisnode プロパティ

analysisnode プロパティ	データ型	プロパティの説明
output_mode	Screen File	出力ノードから生成される出力の、出力先を指定します。
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	string	use_output_name が真 (true) のときに、使用する名前を指定します。
output_format	Text (.txt) HTML (.html) Output (.cou)	出力のタイプを指定します。
by_fields	list	

表 229. analysisnode プロパティ (続き)

analysisnode プロパティ	データ型	プロパティの説明
full_filename	string	ディスク、データ、または HTML の出力を選択した場合の、出力ファイルの名前。
coincidence	フラグ	
performance	フラグ	
evaluation_binary	フラグ	
confidence	フラグ	
threshold	number	
improve_accuracy	number	
field_detection_method	Metadata Name	予測フィールドが元の対象フィールドにどのように一致するかを指定します。Metadata または Name を指定してください。
inc_user_measure	フラグ	
user_if	廃止	
user_then	廃止	
user_else	廃止	
user_compute	[Mean Sum Min Max SDev]	

dataauditnode プロパティ



データ検査ノードでは、欠損値、外れ値、および極値に関する情報の他、各フィールドの要約統計量、ヒストグラムや棒グラフを含む、データを広範に検査するための手段を提供しています。結果は、読みやすいマトリックスで表示されます。ソートして、フルサイズのグラフやデータ準備ノードを生成するために使用できます。

例

```

filenode = stream.createAt("variablefile", "File", 100, 100)
filenode.setPropertyValue("full_filename", "$CLEO_DEMOS/DRUG1n")
node = stream.createAt("dataaudit", "My node", 196, 100)
stream.link(filenode, node)
node.setPropertyValue("custom_fields", True)
node.setPropertyValue("fields", ["Age", "Na", "K"])
node.setPropertyValue("display_graphs", True)
node.setPropertyValue("basic_stats", True)
node.setPropertyValue("advanced_stats", True)
node.setPropertyValue("median_stats", False)
node.setPropertyValue("calculate", ["Count", "Breakdown"])
node.setPropertyValue("outlier_detection_method", "std")
node.setPropertyValue("outlier_detection_std_outlier", 1.0)
node.setPropertyValue("outlier_detection_std_extreme", 3.0)
node.setPropertyValue("output_mode", "Screen")

```


表 230. dataauditnode プロパティ

dataauditnode プロパティ	データ・タイプ	プロパティの説明
custom_fields	フラグ	
fields	[field1 ... fieldN]	
overlay	フィールド	
display_graphs	flag	出力行列中のグラフ表示をオンまたはオフにするために使用されます。
basic_stats	フラグ	
advanced_stats	フラグ	
median_stats	フラグ	
calculate	Count Breakdown	欠損値の計算に使用します。計算方法のいずれか、または両方を選択するか、またはどちらも選択しません。
outlier_detection_method	std iqr	外れ値および極値の検出方法を指定します。
outlier_detection_std_outlier	number	outlier_detection_method が std の場合、外れ値の定義に使用する数値を指定します。
outlier_detection_std_extreme	number	outlier_detection_method が std の場合、外れ値の定義に使用する数値を指定します。
outlier_detection_iqr_outlier	number	outlier_detection_method が iqr の場合、外れ値の定義に使用する数値を指定します。
outlier_detection_iqr_extreme	number	outlier_detection_method が iqr の場合、外れ値の定義に使用する数値を指定します。
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	string	use_output_name が真 (true) のときに、使用する名前を指定します。
output_mode	Screen File	出力ノードから生成される出力の、出力先を指定します。

表 230. dataauditnode プロパティ (続き)

dataauditnode プロパティ	データ・タイプ	プロパティの説明
output_format	Formatted (.tab) Delimited (.csv) HTML (.html) Output (.cou)	出力のタイプを指定します。
paginate_output	フラグ	output_format が HTML の場合、出力がページに分割されるようにします。
lines_per_page	number	paginate_output と共に使用する場合は、出力ページあたりの行数を指定します。
full_filename	string	

extensionoutputnode プロパティ



拡張の出力ノードでは、独自のカスタム R スクリプトまたは Python for Spark スクリプトを使用して、データおよびモデル・スコアリングの結果を分析できます。分析はテキストまたはグラフィックで出力できます。出力はマネージャー領域の「出力」タブに追加されます。あるいは、出力をファイルにリダイレクトできます。

Python for Spark の例

```
##### script example for Python for Spark
import modeler.api
stream = modeler.script.stream()
node = stream.create("extension_output", "extension_output")
node.setPropertyValue("syntax_type", "Python")

python_script = """
import json
import spss.pyspark.runtime

cxt = spss.pyspark.runtime.getContext()
df = cxt.getSparkInputData()
schema = df.dtypes[:]
print df
"""

node.setPropertyValue("python_syntax", python_script)
```

R の例

```
##### script example for R
node.setPropertyValue("syntax_type", "R")
node.setPropertyValue("r_syntax", "print(modelerData$Age)")
```

表 231. extensionoutputnode プロパティ

extensionoutputnode プロパティ	データ・タイプ	プロパティの説明
syntax_type	R Python	R または Python のどちらのスクリプトを実行するか指定します (R がデフォルトです)。
r_syntax	string	モデル・スコアリング用の R スクリプト・シンタックス。
python_syntax	string	モデル・スコアリング用の Python スクリプト・シンタックス。
convert_flags	StringsAndDoubles LogicalValues	フラグ型フィールドを変換するためのオプション。
convert_missing	フラグ	欠損値を R の NA 値に変換するオプションです。
convert_datetime	フラグ	日付形式または日付/時刻形式の変数を R の日付/時刻形式に変換するためのオプション。
convert_datetime_class	POSIXct POSIXlt	日付形式または日付/時刻形式の変数のうち、どの形式の変数を変換するかを指定するためのオプション。
output_to	Screen File	出力形式 (Screen または File) を指定します。
output_type	Graph Text	グラフィカル出力またはテキスト出力のどちらを作成するか指定します。
full_filename	string	生成された出力に使用するファイル名。
graph_file_type	HTML COU	出力ファイルのファイルの種類 (.html または .cou)。
text_file_type	HTML TEXT COU	テキスト出力のファイルの種類 (.html、.txt、または .cou) を指定します。

kdeexport プロパティ



カーネル密度推定 (KDE)[®] は、Ball Tree または KD Tree のアルゴリズムを使用してクエリを効率化し、教師なし学習、特徴量エンジニアリング、データのモデル化の概念を結合します。KDE などの近隣ベースの手法が、最もよく使用され、有用な密度推定手法です。スボス モデラー の KDE モデル作成および KDE シミュレーションのノードは、KDE ライブラリーのコア機能およびよく使用されるパラメータを公開します。これらのノードは Python で実装されています。

表 232. kdeexport プロパティ

kdeexport プロパティ	データ・タイプ	プロパティの説明
bandwidth	DOUBLE	デフォルトは 1 です。
kernel	string	使用するカーネル: gaussian または tophat。デフォルトは gaussian です。
algorithm	string	使用するツリー・アルゴリズム: kd_tree、ball_tree、または auto。デフォルトは auto です。
metric	string	距離を計算するときに使用するメトリック。kd_tree アルゴリズムの場合は、Euclidean、Chebyshev、Cityblock、Minkowski、Manhattan、Infinity、P、L2、または L1 から選択します。ball_tree アルゴリズムの場合は、Euclidian、Braycurtis、Chebyshev、Canberra、Cityblock、Dice、Hamming、Infinity、Jaccard、L1、L2、Minkowski、Matching、Manhattan、P、Rogersanimoto、Russellrao、Sokalmichener、Sokalsneath、または Kulsinski から選択します。デフォルトは Euclidean です。
atol	Float	希望する結果の絶対許容度。一般に、許容度を大きくすると、実行速度が上がります。デフォルトは 0.0 です。
rtol	Float	希望する結果の相対許容度。一般に、許容度を大きくすると、実行速度が上がります。デフォルトは 1E-8 です。
breadthFirst	Boolean	幅優先のアプローチを使用するには、True に設定します。深さ優先のアプローチを使用するには、False に設定します。デフォルトは True です。
LeafSize	integer	基本となるツリーのリーフ サイズ。デフォルトは 40 です。この値を変更すると、パフォーマンスに大きな影響を与える可能性があります。
pValue	DOUBLE	メトリックに Minkowski を使用している場合は、使用する P 値を指定します。デフォルトは 1.5 です。

matrixnode プロパティ



クロス集計ノードで、フィールド間の関係を示すテーブルを作成します。これは、通常、2つのシンボル値フィールド間の関係を示すために使用されますが、フラグ型フィールドまたは数値フィールド間の関係を示すこともできます。

例

```
node = stream.create("matrix", "My node")
# "Settings" tab
node.setPropertyValue("fields", "Numerics")
node.setPropertyValue("row", "K")
node.setPropertyValue("column", "Na")
node.setPropertyValue("cell_contents", "Function")
node.setPropertyValue("function_field", "Age")
node.setPropertyValue("function", "Sum")
# "Appearance" tab
node.setPropertyValue("sort_mode", "Ascending")
node.setPropertyValue("highlight_top", 1)
node.setPropertyValue("highlight_bottom", 5)
node.setPropertyValue("display", ["Counts", "Expected", "Residuals"])
node.setPropertyValue("include_totals", True)
# "Output" tab
node.setPropertyValue("full_filename", "C:/output/matrix_output.html")
node.setPropertyValue("output_format", "HTML")
node.setPropertyValue("paginate_output", True)
node.setPropertyValue("lines_per_page", 50)
```

表 233. *matrixnode* プロパティ

matrixnode プロパティ	データ型	プロパティの説明
fields	Selected Flags Numerics	
row	フィールド	
column	フィールド	
include_missing_values	フラグ	ユーザーによる欠損値 (空白) とシステムによる欠損値 (ヌル) が、行と列の出力に含まれるかどうかを指定します。
cell_contents	CrossTabs Function	
function_field	string	
function	Sum Mean Min Max SDev	

表 233. *matrixnode* プロパティ (続き)

matrixnode プロパティ	データ型	プロパティの説明
sort_mode	Unsorted Ascending Descending	
highlight_top	<i>number</i>	ゼロでない場合に真 (true)。
highlight_bottom	<i>number</i>	ゼロでない場合に真 (true)。
display	[Counts Expected Residuals RowPct ColumnPct TotalPct]	
include_totals	フラグ	
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	<i>string</i>	<code>use_output_name</code> が真 (true) のときに、使用する名前を指定します。
output_mode	Screen File	出力ノードから生成される出力の、出力先を指定します。
output_format	Formatted (.tab) Delimited (.csv) HTML (.html) Output (.cou)	出力のタイプを指定します。 Formatted と Delimited の両方が、テーブル内で行と列を入れ替える修飾子 transposed を伴うことができます。
paginate_output	フラグ	<code>output_format</code> が HTML の場合、出力がページに分割されるようにします。
lines_per_page	<i>number</i>	<code>paginate_output</code> と共に使用する場合は、出力ページあたりの行数を指定します。
full_filename	<i>string</i>	

meansnode プロパティ



平均比較ノードでは、独立したグループ間で、または関連するフィールドのペア間で著しい違いがあるかどうかを調べるために、平均を比較します。例えば、販売促進活動の前後で平均収益を比較したり、販売促進活動を受けなかった顧客と受けた顧客からの収益を比較することができます。

例

```
node = stream.create("means", "My node")
node.setPropertyValue("means_mode", "BetweenFields")
node.setPropertyValue("paired_fields", [["OPEN_BAL", "CURR_BAL"]])
node.setPropertyValue("label_correlations", True)
node.setPropertyValue("output_view", "Advanced")
node.setPropertyValue("output_mode", "File")
node.setPropertyValue("output_format", "HTML")
node.setPropertyValue("full_filename", "C:/output/means_output.html")
```

表 234. meansnode プロパティ

meansnode プロパティ	データ・タイプ	プロパティの説明
means_mode	BetweenGroups BetweenFields	データに実行する平均統計処理の種類を指定します。
test_fields	[field1 ... fieldn]	means_mode が BetweenGroups に設定されているときのテスト・フィールドを指定します。
grouping_field	フィールド	グループにまとめるフィールドを指定します。
paired_fields	[[field1 field2] [field3 field4] ...]	means_mode が BetweenFields に設定されているときに使用するフィールドのペアを指定します。
label_correlations	フラグ	相関ラベルが出力に表示されるかどうかを指定します。 means_mode が BetweenFields に設定されているときにのみ、この設定が適用されます。
correlation_mode	Probability Absolute	確率 (Probability) または絶対値 (Absolute) のどちらかで相関にラベルを付けることを指定します。
weak_label	string	
medium_label	string	
strong_label	string	

表 234. meansnode プロパティ (続き)

meansnode プロパティ	データ・タイプ	プロパティの説明
weak_below_probability	number	correlation_mode が Probability に設定されているときに、弱い相関の分割値を指定します。この値は、例えば 0.90 のように、0 と 1 の間にする必要があります。
strong_above_probability	number	強い相関の分割値。
weak_below_absolute	number	correlation_mode が Absolute に設定されているときに、弱い相関の分割値を指定します。この値は、例えば 0.90 のように、0 と 1 の間にする必要があります。
strong_above_absolute	number	強い相関の分割値。
unimportant_label	string	
marginal_label	string	
important_label	string	
unimportant_below	number	低いフィールド重要度の分割値。この値は、例えば 0.90 のように、0 と 1 の間にする必要があります。
important_above	number	
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	string	使用する名前。
output_mode	Screen File	出力ノードから生成された出力の出力先を指定します。
output_format	Formatted (.tab) Delimited (.csv) HTML (.html) Output (.cou)	出力のタイプを指定します。
full_filename	string	
output_view	Simple Advanced	出力に単純な (Simple) ビューが表示されるか、または詳細な (Advanced) ビューが表示されるかを指定します。

reportnode プロパティ



レポート・ノードで、固定テキスト、およびデータやデータから導かれた他の式を含む、フォーマット済みレポートを作成します。レポートの書式は、固定テキストとデータの出力構成を定義するテキスト テンプレートを使用して指定します。テンプレート内の HTML タグを使用し、また「出力」タブでオプションを設定することで、カスタムのテキスト書式設定を提供できます。テンプレートに CLEM 式を使用して、データ値およびその他の条件付き出力を含めることができます。

例

```
node = stream.create("report", "My node")
node.setPropertyValue("output_format", "HTML")
node.setPropertyValue("full_filename", "C:/report_output.html")
node.setPropertyValue("lines_per_page", 50)
node.setPropertyValue("title", "Report node created by a script")
node.setPropertyValue("highlights", False)
```

表 235. reportnode プロパティ

reportnode プロパティ	データ・タイプ	プロパティの説明
output_mode	Screen File	出力ノードから生成される出力の、出力先を指定します。
output_format	HTML (.html) Text (.txt) Output (.cou)	ファイル出力のタイプを指定します。
format	Auto Custom	出力を自動的にフォーマット設定するか、テンプレートに含まれる HTML を使用してフォーマット設定するかを選択するために使用します。テンプレート内の HTML フォーマット設定を使用するには、Custom を指定します。
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	string	use_output_name が真 (true) のときに、使用する名前を指定します。
text	string	
full_filename	string	
highlights	フラグ	
title	string	
lines_per_page	number	

rouputnode のプロパティ



R 出力ノードでは、独自のカスタム R スクリプトを使用して、データおよびモデル・スコアリングの結果を分析できます。分析はテキストまたはグラフィックで出力できます。出力はマネージャー領域の「出力」タブに追加されます。あるいは、出力をファイルにリダイレクトできます。

表 236. rouputnode のプロパティ		
rouputnode のプロパティ	データ・タイプ	プロパティの説明
syntax	string	
convert_flags	StringsAndDoubles LogicalValues	
convert_datetime	フラグ	
convert_datetime_class	POSIXct POSIXlt	
convert_missing	フラグ	
output_name	Auto Custom	
custom_name	string	
output_to	Screen File	
output_type	Graph Text	
full_filename	string	
graph_file_type	HTML COU	
text_file_type	HTML TEXT COU	

setglobalsnode プロパティ



グローバルの設定ノードで、データを走査し、CLEM 式で利用できる要約値を算出します。例えば、このノードを使用して年齢というフィールドの統計量を計算し、関数 @GLOBAL_MEAN(age) を挿入することにより、CLEM 式で年齢の全体平均値を使用することができます。

例

```
node = stream.create("setglobals", "My node")
node.setKeyedPropertyValue("globals", "Na", ["Max", "Sum", "Mean"])
node.setKeyedPropertyValue("globals", "K", ["Max", "Sum", "Mean"])
```

```
node.setKeyedPropertyValue("globals", "Age", ["Max", "Sum", "Mean", "SDev"])
node.setPropertyValue("clear_first", False)
node.setPropertyValue("show_preview", True)
```

表 237. *setglobalsnode* プロパティ

setglobalsnode プロパティ	データ・タイプ	プロパティの説明
globals	[Sum Mean Min Max SDev]	フィールドを設定する構造化プロパティは、次の形式で参照する必要があります。 node.setKeyedPropertyValue("globals", "Age", ["Max", "Sum", "Mean", "SDev"])
clear_first	フラグ	
show_preview	フラグ	

simevalnode プロパティ



シミュレーション評価ノードは、指定された予測される対象フィールドを評価し、対象フィールドの分布と相関情報を提供します。

表 238. *simevalnode* プロパティ

simevalnode プロパティ	データ・タイプ	プロパティの説明
target	フィールド	
iteration	フィールド	
presorted_by_iteration	Boolean	
max_iterations	number	
tornado_fields	[field1...fieldN]	
plot_pdf	Boolean	
plot_cdf	Boolean	
show_ref_mean	Boolean	
show_ref_median	Boolean	
show_ref_sigma	Boolean	
num_ref_sigma	number	
show_ref_pct	Boolean	
ref_pct_bottom	number	
ref_pct_top	number	
show_ref_custom	Boolean	
ref_custom_values	[number1...numberN]	

表 238. *simevalnode* プロパティ (続き)

simevalnode プロパティ	データ・タイプ	プロパティの説明
category_values	Category Probabilities Both	
category_groups	Categories Iterations	
create_pct_table	<i>Boolean</i>	
pct_table	Quartiles Intervals Custom	
pct_intervals_num	<i>number</i>	
pct_custom_values	[<i>number1...numberN</i>]	

simfitnode プロパティ



シミュレーションの当てはめノードは、各フィールドのデータの統計的な分布を調べ、最も適合する分布を各フィールドに割り当ててシミュレーション生成ノードを生成 (または更新) します。この後、シミュレーション生成ノードを使用して、シミュレートするデータを生成することができます。

表 239. *simfitnode* プロパティ

simfitnode プロパティ	データ型	プロパティの説明
build	Node XMLExport Both	
use_source_node_name	<i>Boolean</i>	
source_node_name	<i>string</i>	生成または更新される入力ノードのカスタム名。
use_cases	All LimitFirstN	
use_case_limit	<i>integer</i>	
fit_criterion	AndersonDarling KolmogorovSmirnov	
num_bins	<i>integer</i>	
parameter_xml_filename	<i>string</i>	
generate_parameter_import	<i>Boolean</i>	

statisticsnode プロパティ



記述統計ノードでは、数値型フィールドに関する基本的な集計情報が提供されます。このノードで、個々のフィールドの要約統計量とフィールド間の相関が計算されます。

例

```
node = stream.create("statistics", "My node")
# "Settings" tab
node.setPropertyValue("examine", ["Age", "BP", "Drug"])
node.setPropertyValue("statistics", ["mean", "sum", "sdev"])
node.setPropertyValue("correlate", ["BP", "Drug"])
# "Correlation Labels..." section
node.setPropertyValue("label_correlations", True)
node.setPropertyValue("weak_below_absolute", 0.25)
node.setPropertyValue("weak_label", "lower quartile")
node.setPropertyValue("strong_above_absolute", 0.75)
node.setPropertyValue("medium_label", "middle quartiles")
node.setPropertyValue("strong_label", "upper quartile")
# "Output" tab
node.setPropertyValue("full_filename", "c:/output/statistics_output.html")
node.setPropertyValue("output_format", "HTML")
```

表 240. statisticsnode プロパティ

statisticsnode プロパティ	データ・タイプ	プロパティの説明
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	string	use_output_name が真 (true) のときに、使用する名前を指定します。
output_mode	Screen File	出力ノードから生成される出力の、出力先を指定します。
output_format	Text (.txt) HTML (.html) Output (.cou)	出力のタイプを指定します。
full_filename	string	
examine	リスト	
correlate	リスト	
statistics	[count mean sum min max range variance sdev semean median mode]	

表 240. statisticsnode プロパティ (続き)

statisticsnode プロパティ	データ・タイプ	プロパティの説明
correlation_mode	Probability Absolute	確率 (Probability) または絶対値 (Absolute) のどちらかで相関にラベルを付けることを指定します。
label_correlations	フラグ	
weak_label	string	
medium_label	string	
strong_label	string	
weak_below_probability	number	correlation_mode が Probability に設定されているときに、弱い相関の分割値を指定します。この値は、例えば 0.90 のように、0 と 1 の間にする必要があります。
strong_above_probability	number	強い相関の分割値。
weak_below_absolute	number	correlation_mode が Absolute に設定されているときに、弱い相関の分割値を指定します。この値は、例えば 0.90 のように、0 と 1 の間にする必要があります。
strong_above_absolute	number	強い相関の分割値。

statisticsoutputnode プロパティ



Statistics 出力ノードを使用すると、IBM SPSS 統計 手続きを呼び出し、IBM スポス モデラー データを分析することができます。さまざまな IBM SPSS 統計 分析手続き にアクセスできます。このノードには、ライセンス交付を受けた IBM SPSS 統計のコピーが必要です。

このノードのプロパティについては、[425 ページ](#)の『[statisticsoutputnode プロパティ](#)』に記載されています。

tablenode プロパティ



テーブル・ノードで、データがテーブル形式で表示されます。このデータは、ファイルにも書き込めます。この機能は、データの値を調査したり、データを読みやすい形式でエクスポートする必要がある場合に役立ちます。

例

```
node = stream.create("table", "My node")
node.setPropertyValue("highlight_expr", "Age > 30")
node.setPropertyValue("output_format", "HTML")
node.setPropertyValue("transpose_data", True)
node.setPropertyValue("full_filename", "C:/output/table_output.htm")
```

```
node.setPropertyValue("paginate_output", True)
node.setPropertyValue("lines_per_page", 50)
```

表 241. *tablenode* プロパティ

tablenode プロパティ	データ・タイプ	プロパティの説明
full_filename	<i>string</i>	ディスク、データ、または HTML の出力を選択した場合の、出力ファイルの名前。
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	<i>string</i>	<i>use_output_name</i> が真 (true) のときに、使用する名前を指定します。
output_mode	Screen File	出力ノードから生成される出力の、出力先を指定します。
output_format	Formatted (.tab) Delimited (.csv) HTML (.html) Output (.cou)	出力のタイプを指定します。
transpose_data	フラグ	エクスポート前にデータの行列を入れ替えて、行がフィールドを、列がレコードを表すようにします。
paginate_output	フラグ	<i>output_format</i> が HTML の場合、出力がページに分割されるようにします。
lines_per_page	<i>number</i>	<i>paginate_output</i> と共に使用する場合は、出力ページあたりの行数を指定します。
highlight_expr	<i>string</i>	
output	<i>string</i>	ノードで直前に構築されたテーブルへの参照を保持する、読み取り専用プロパティ。
value_labels	[[<i>Value LabelString</i>] [<i>Value LabelString</i>] ...]	値のペアのためのラベルを指定します。
display_places	<i>integer</i>	フィールドが表示されるときの小数部の桁数を設定します (REAL ストレージのフィールドにのみ適用)。-1 を設定すると、ストリームのデフォルトが使用されます。

表 241. *tablenode* プロパティ (続き)

tablenode プロパティ	データ・タイプ	プロパティの説明
export_places	<i>integer</i>	フィールドが出力されるときの小数部の桁数を設定します (REAL ストレージのフィールドにのみ適用)。-1 を設定すると、ストリームのデフォルトが使用されます。
decimal_separator	DEFAULT PERIOD COMMA	フィールドの小数点記号を指定します (REAL ストレージのフィールドにのみ適用)。
date_format	"DDMMYY" "MMDDYY" "YYMMDD" "YYYYMMDD" "YYYYDDD" DAY MONTH "DD-MM-YY" "DD-MM-YYYY" "MM-DD-YY" "MM-DD-YYYY" "DD-MON-YY" "DD-MON-YYYY" "YYYY-MM-DD" "DD.MM.YY" "DD.MM.YYYY" "MM.DD.YYYY" "DD.MON.YY" "DD.MON.YYYY" "DD/MM/YY" "DD/MM/YYYY" "MM/DD/YY" "MM/DD/YYYY" "DD/MON/YY" "DD/MON/YYYY" MON YYYY q Q YYYY ww WK YYYY	フィールドの日付形式を設定します (DATE または TIMESTAMP ストレージのフィールドにのみ適用されます)。

表 241. *tablenode* プロパティ (続き)

tablenode プロパティ	データ・タイプ	プロパティの説明
time_format	"HHMMSS" "HHMM" "MMSS" "HH:MM:SS" "HH:MM" "MM:SS" "(H)H:(M)M:(S)S" "(H)H:(M)M" "(M)M:(S)S" "HH.MM.SS" "HH.MM" "MM.SS" "(H)H.(M)M.(S)S" "(H)H.(M)M" "(M)M.(S)S"	フィールドの日付形式を設定します (TIME または TIMESTAMP ストレージのフィールドにのみ適用されます)。
column_width	<i>integer</i>	フィールドに列幅を設定します。-1 という値を指定すると、列幅は Auto に設定されます。
justify	AUTO CENTER LEFT RIGHT	フィールドに列調整を設定します。

transformnode プロパティ



変換ノードを使用すると、選択したフィールドに変換の結果を適用する前に、それらの結果を選択して視覚的にプレビューすることができます。

例

```
node = stream.create("transform", "My node")
node.setPropertyValue("fields", ["AGE", "INCOME"])
node.setPropertyValue("formula", "Select")
node.setPropertyValue("formula_log_n", True)
node.setPropertyValue("formula_log_n_offset", 1)
```

表 242. *transformnode* プロパティ

transformnode プロパティ	データ・タイプ	プロパティの説明
fields	[<i>field1</i> ... <i>fieldn</i>]	変換で使用するフィールド。
formula	All Select	すべての変換を計算するか、選択した変換を計算するかを指定します。
formula_inverse	フラグ	逆変換を使用するかどうかを指定します。
formula_inverse_offset	<i>number</i>	式で使用するデータ・オフセットを指定します。ユーザーが指定しない限り、デフォルトで 0 に設定されます。
formula_log_n	フラグ	$\log_n$ 変換を使用するかどうかを指定します。
formula_log_n_offset	<i>number</i>	
formula_log_10	フラグ	$\log_{10}$ 変換を使用するかどうかを指定します。
formula_log_10_offset	<i>number</i>	
formula_exponential	フラグ	指数変換 (e^x) を使用するかどうかを指定します。
formula_square_root	フラグ	平方根変換を使用するかどうかを指定します。
use_output_name	フラグ	ユーザー設定の出力名が使用されるかどうかを指定します。
output_name	<i>string</i>	<code>use_output_name</code> が真 (true) のときに、使用する名前を指定します。
output_mode	Screen File	出力ノードから生成される出力の、出力先を指定します。
output_format	HTML (.html) Output (.cou)	出力のタイプを指定します。
paginate_output	フラグ	<code>output_format</code> が HTML の場合、出力がページに分割されるようにします。

表 242. <i>transformnode</i> プロパティ (続き)		
transformnode プロパティ	データ・タイプ	プロパティの説明
lines_per_page	<i>number</i>	paginate_output と共に使用する場合は、出力ページあたりの行数を指定します。
full_filename	<i>string</i>	ファイル出力に使用するファイル名を指定します。

第 17 章 エクスポート・ノードのプロパティ

共通のエクスポート・ノード・プロパティ

次のプロパティは、すべてのエクスポート・ノードに共通しています。

表 243. 共通のエクスポート・ノード・プロパティ		
プロパティ	値	プロパティの説明
publish_path	string	公開されたイメージおよびパラメーター・ファイルに使用するルート名を指定します。
publish_metadata	フラグ	イメージの入力および出力、それらのデータ・モデルを説明するメタデータ・ファイルを作成するかどうかを指定します。
publish_use_parameters	フラグ	ストリーム・パラメーターが*.par ファイルに含まれるかどうかを指定します。
publish_parameters	文字列のリスト	使用するパラメーターを指定します。
execute_mode	export_data publish	ストリームを公開せずにノードを実行するかどうか、ノードの実行時にストリームを自動的に公開するかどうかを指定します。

asexport プロパティ

分析サーバー エクスポートにより、Hadoop 分散ファイル・システム (HDFS) でストリームを実行することができます。

例

```
node.setPropertyValue("use_default_as", False)
node.setPropertyValue("connection",
["false", "9.119.141.141", "9080", "analyticserver", "ibm", "admin", "admin", "false", "", "", "", "", ""])
```

表 244. asexport プロパティ		
asexport プロパティ	データ・タイプ	プロパティの説明
data_source	string	データ・ソースの名前。
export_mode	string	エクスポートしたデータを既存のデータ・ソースに追加する (append) か、既存のデータ・ソースを上書きする (overwrite) かを指定します。

表 244. asexport プロパティ (続き)

asexport プロパティ	データ・タイプ	プロパティの説明
use_default_as	Boolean	True に設定すると、サーバーの options.cfg ファイルで構成されているデフォルトの分析サーバー 接続が使用されます。False に設定した場合、このノードの接続が使用されます。
connection	["string", "string", "string", "string", "string", "string", "string", "string", "string", "string", "string", "string", "string"]	分析サーバー 接続の詳細を含むリストのプロパティ。形式は ["is_secure_connect", "server_url", "server_port", "context_root", "consumer", "user_name", "password", "use-kerberos-auth", "kerberos-krb5-config-file-path", "kerberos-jaas-config-file-path", "kerberos-krb5-service-principal-name", "enable-kerberos-debug"] です。ここで、is_secure_connect: はセキュア接続が使用されているかどうかを示し、true または false のいずれかです。use-kerberos-auth: は、Kerberos 認証が使用されるかどうかを示します。これは、true または false のいずれかです。enable-kerberos-debug: は、Kerberos 認証のデバッグ・モードが使用されているかどうかを示します。これは、true または false のいずれかです。

cognosexportnode プロパティ



IBM Cognos エクスポート・ノードは、Cognos データベースで読み取ることができる形式でデータをエクスポートできます。

このノードの場合、Cognos 接続と ODBC 接続を定義する必要があります。

Cognos 接続

Cognos 接続のプロパティは次のとおりです。

表 245. *cognosexportnode* プロパティ

cognosexportnode プロパティ	データ・タイプ	プロパティの説明
cognos_connection	["文字列","フラグ","文字列","文字列","文字列"]	Cognos サーバーの接続の詳細を含むリストのプロパティ。形式は以下のとおりです。 <pre>["Cognos_server_URL", login_mode, "namespace", "username", "password"]</pre> ここで、 Cognos_server_URL は、ソースが格納されている Cognos サーバーの URL です。 login_mode は、匿名ログインが使用されるかどうかを示します。true または false のいずれかです。true に設定されている場合、以下のフィールドを "." に設定する必要があります。 namespace はサーバーへのログオンに使用するセキュリティー認証プロバイダを示します。 username および password は Cognos サーバーにログオンする際に使用するユーザー名とパスワードです。 login_mode の代わりに、以下のモードも使用可能です。 - anonymousMode 以下に例を示します。['Cognos_server_url', 'anonymousMode', "namespace", "username", "password"] - credentialMode 以下に例を示します。['Cognos_server_url', 'credentialMode', "namespace", "username", "password"]

表 245. *cognosexportnode* プロパティ (続き)

cognosexportnode プロパティ	データ・タイプ	プロパティの説明
		storedCredentialMode 以下に例を示します。 ['Cognos_server_url', 'storedCredentialMode', "stored_credential_name"] ここで、stored_credential_name は、リポジトリ内での Cognos の資格情報の名前です。
cognos_package_name	string	データをエクスポートしている Cognos データ・ソース (通常はデータベース) のパスおよび名前。次に例を示します。 /Public Folders/MyPackage
cognos_datasource	string	
cognos_export_mode	Publish ExportFile	
cognos_filename	string	

ODBC 接続

ODBC 接続のプロパティは次のセクションの *databaseexportnode* に示されているものと同じです。ただし、datasource プロパティは有効ではありません。

databaseexportnode プロパティ



データベース・エクスポート・ノードで、データを ODBC 対応のリレーショナル・データ・ソースに書き込みます。ODBC データ・ソースに書き込むには、データ・ソースが存在し、それに対する書き込み権限を持っている必要があります。

例

```

...
Assumes a datasource named "MyDatasource" has been configured
...
stream = modeler.script.stream()
db_exportnode = stream.createAt("databaseexport", "DB Export", 200, 200)
applynn = stream.findByType("applyneuralnetwork", None)
stream.link(applynn, db_exportnode)

# Export tab
db_exportnode.setPropertyValue("username", "user")
db_exportnode.setPropertyValue("datasource", "MyDatasource")
db_exportnode.setPropertyValue("password", "password")
db_exportnode.setPropertyValue("table_name", "predictions")
db_exportnode.setPropertyValue("write_mode", "Create")

```



```

db_exportnode.setPropertyValue("generate_import", True)
db_exportnode.setPropertyValue("drop_existing_table", True)
db_exportnode.setPropertyValue("delete_existing_rows", True)
db_exportnode.setPropertyValue("default_string_size", 32)

# Schema dialog
db_exportnode.setKeyedPropertyValue("type", "region", "VARCHAR(10)")
db_exportnode.setKeyedPropertyValue("export_db_primarykey", "id", True)
db_exportnode.setPropertyValue("use_custom_create_table_command", True)
db_exportnode.setPropertyValue("custom_create_table_command", "My SQL Code")

# Indexes dialog
db_exportnode.setPropertyValue("use_custom_create_index_command", True)
db_exportnode.setPropertyValue("custom_create_index_command", "CREATE BITMAP
INDEX <index-name>
ON <table-name> <(index-columns)>")
db_exportnode.setKeyedPropertyValue("indexes", "MYINDEX", ["fields", ["id",
"region"]])

```

表 246. databaseexportnode プロパティ

databaseexportnode プロパティ	データ・タイプ	プロパティの説明
datasource	string	
username	string	
password	string	
epassword	string	このスロットは、実行時に読み込み用になります。暗号化パスワードを生成するには、「ツール」メニューの「パスワード暗号化ツール」を使用してください。詳しくは、トピック 54 ページの『 暗号化パスワードの生成 』を参照してください。
table_name	string	
write_mode	Create Append Merge	
map	string	ストリーム・フィールド名をデータベース列名にマッピングします (write_mode が Merge の場合にのみ有効)。 結合の場合、すべてのフィールドをマッピングしてエクスポートする必要があります。データベース内に存在しないフィールド名が、新しい列として追加されます。

表 246. databaseexportnode プロパティ (続き)

databaseexportnode プロパティ	データ・タイプ	プロパティの説明
key_fields	リスト	キーに使用されるストリーム・フィールドを指定します。map プロパティは、データベースでストリーム・フィールド内で対応する内容を表示します。
join	Database Add	
drop_existing_table	フラグ	
delete_existing_rows	フラグ	
default_string_size	integer	
type		スキーマタイプの設定に用いられる構造化プロパティ。
generate_import	フラグ	
use_custom_create_table_command	フラグ	custom_create_table スロットを使用して、標準の CREATE TABLE SQL コマンドを変更します。
custom_create_table_command	string	標準の CREATE TABLE SQL コマンドの代わりに使用する文字列コマンドを指定します。
use_batch	フラグ	次のプロパティは、データベースのバルク・ロード用の詳細オプションです。Use_batch の値が True の場合は、データベースへの行ごとのコミットがオフになります。
batch_size	number	メモリーにコミットする前にデータベースに送信するレコード数を指定します。
bulk_loading	Off ODBC External	バルク・ロードの種類を指定します。ODBC および External 用の付加オプションを次に示します。
not_logged	フラグ	
odbc_binding	Row Column	ODBC 経由のバルク・ロードにおける、行方向または列方向のバインドを指定します。

表 246. databaseexportnode プロパティ (続き)

databaseexportnode プロパティ	データ・タイプ	プロパティの説明
loader_delimit_mode	Tab Space Other	外部プログラム経由のバルク・ロードの場合に、区切り文字の種類を指定します。Other は、 loader_other_delimiter プロパティと組み合わせて選択し、コンマ (,) のような区切り文字を指定します。
loader_other_delimiter	string	
specify_data_file	フラグ	真 (True) のフラグを設定すると、以下の data_file プロパティが有効になります。このプロパティには、データベースにバルクロードする際の書き込み先のファイル名とパスを指定することができます。
data_file	string	
specify_loader_program	フラグ	真 (True) のフラグを設定すると、以下の loader_program プロパティが有効になります。このプロパティには、外部ローダスクリプトまたはプログラムの名前と場所を指定することができます。
loader_program	string	
gen_logfile	フラグ	真 (True) のフラグを設定すると、以下の logfile_name が有効になります。このプロパティには、エラー ログを生成するための、サーバー上のファイル名を指定することができます。
logfile_name	string	
check_table_size	フラグ	真 (True) のフラグを設定すると、IBM スポス モデラー からエクスポートされる行数に対応してデータベースのテーブル サイズを確実に増加させるために、テーブル検査が実施されます。
loader_options	string	ローダー・プログラムに対して、-comment および -specialdir のような、他の引数を指定します。
export_db_primarykey	フラグ	指定されたフィールドがプライマリ キーかどうかを指定します。

表 246. *databaseexportnode* プロパティ (続き)

databaseexportnode プロパティ	データ・タイプ	プロパティの説明
<code>use_custom_create_index_command</code>	フラグ	<code>true</code> の場合、すべてのインデックスに対してカスタム SQL (ユーザー指定の SQL) を有効にします。
<code>custom_create_index_command</code>	<i>string</i>	カスタム SQL (ユーザー指定の SQL) が有効にされている場合、インデックスの作成に使用される SQL コマンドを指定します。(この値は、下に示す特定のインデックスに対して上書きできます。)
<code>indexes.INDEXNAME.fields</code>		必要の場合は指定されたインデックスを作成し、そのインデックスに含まれるフィールド名を一覧表示します。
INDEXNAME "use_custom_create_index_command"	フラグ	特定のインデックスに対してカスタム SQL (ユーザー指定の SQL) を有効または無効にするのに使用されます。後続の表の後にある例を参照してください。
INDEXNAME "custom_create_index_command"	<i>string</i>	指定されたインデックスに使用されるカスタム SQL (ユーザー指定の SQL) を使用します。後続の表の後にある例を参照してください。
<code>indexes.INDEXNAME.remove</code>	フラグ	<code>True</code> の場合、指定されたインデックスをインデックスのセットから削除します。
<code>table_space</code>	<i>string</i>	作成されるテーブル・スペースを指定します。
<code>use_partition</code>	フラグ	分布ハッシュ・フィールドが使用されるよう指定します。
<code>partition_field</code>	<i>string</i>	分布ハッシュ・フィールドの内容を消去します。

注：一部のデータベースでは、データベース表が圧縮を使用してエクスポート用に作成されるように指定できます (例えば、SQL の `CREATE TABLE MYTABLE (...) COMPRESS YES;` に相当します)。次のようにプロパティ `use_compression` および `compression_mode` を指定して、この機能をサポートします。

表 247. 圧縮機能を使用した *databaseexportnode* プロパティ

databaseexportnode プロパティ	データ・タイプ	プロパティの説明
<code>use_compression</code>	<i>Boolean</i>	<code>True</code> に設定した場合は、圧縮によるエクスポート用のテーブルを作成します。

表 247. 圧縮機能を使用した *databaseexportnode* プロパティ (続き)

databaseexportnode プロパティ	データ・タイプ	プロパティの説明
compression_mode	Row	SQL Server データベースの圧縮レベルを設定します。
	Page	
	Default	Oracle データベースの圧縮レベルを設定します。値 OLTP、Query_High、Query_Low、Archive_High、および Archive_Low には最低限 Oracle 11gR2 が必要です。
	Direct_Load_Operations	
	All_Operations	
	Basic	
	OLTP	
	Query_High	
	Query_Low	
	Archive_High	
	Archive_Low	

CREATE INDEX コマンドを特定のインデックス用に変更する方法を示す例:

```
db_exportnode.setKeyedPropertyValue("indexes", "MYINDEX",
["use_custom_create_index_command",
True])db_exportnode.setKeyedPropertyValue("indexes", "MYINDEX",
["custom_create_index_command",
"CREATE BITMAP INDEX <index-name> ON <table-name> <(index-columns)>"])
```

あるいは、同じ処理をハッシュ テーブルを用いて行うこともできます。

```
db_exportnode.setKeyedPropertyValue("indexes", "MYINDEX", ["fields":["id",
"region"],
"use_custom_create_index_command":True,
"custom_create_index_command":"CREATE INDEX <index-name> ON
<table-name> <(index-columns)>"])
```

datacollectionexportnode プロパティ



データ収集 エクスポート・ノードは、データ収集 の市場調査ソフトウェアで使用する形式でデータを出力します。このノードを使用するには、データ収集 データ・ライブラリーをインストールする必要があります。

例

```
stream = modeler.script.stream()
datacollectionexportnode = stream.createAt("datacollectionexport", "Data
Collection", 200, 200)
```

```

datacollectionexportnode.setPropertyValue("metadata_file", "c:\\museums.mdd")
datacollectionexportnode.setPropertyValue("merge_metadata", "Overwrite")
datacollectionexportnode.setPropertyValue("casedata_file", "c:\\
\\museumdata.sav")
datacollectionexportnode.setPropertyValue("generate_import", True)
datacollectionexportnode.setPropertyValue("enable_system_variables", True)

```

表 248. datacollectionexportnode プロパティ

datacollectionexportnode プロパティ	データ・タイプ	プロパティの説明
metadata_file	string	出力するメタデータ・ファイルの名前。
merge_metadata	Overwrite MergeCurrent	
enable_system_variables	フラグ	エクスポートされた .mdd ファイルに データ収集 システム変数を含むかどうかを指定します。
casedata_file	string	ケース・データがエクスポートされる .sav ファイルの名前。
generate_import	フラグ	

excelexportnode プロパティ



Excel エクスポート・ノードでは、Microsoft Excel にデータを出力します。xlsx ファイル・フォーマット。オプションで、ノードが実行されるときに自動的に Excel が起動し、エクスポートするファイルを開けるように選択できます。

例

```

stream = modeler.script.stream()
excelexportnode = stream.createAt("excelexport", "Excel", 200, 200)
excelexportnode.setPropertyValue("full_filename", "C:/output/myexport.xlsx")
excelexportnode.setPropertyValue("excel_file_type", "Excel2007")
excelexportnode.setPropertyValue("inc_field_names", True)
excelexportnode.setPropertyValue("inc_labels_as_cell_notes", False)
excelexportnode.setPropertyValue("launch_application", True)
excelexportnode.setPropertyValue("generate_import", True)

```

表 249. excelexportnode プロパティ

excelexportnode プロパティ	データ型	プロパティの説明
full_filename	string	
excel_file_type	Excel2007	
export_mode	Create Append	
inc_field_names	フラグ	フィールド名がワークシートの最初の行に表示されるかどうかを指定します。

表 249. excelexportnode プロパティ (続き)

excelexportnode プロパティ	データ型	プロパティの説明
start_cell	string	エクスポートの開始セルを指定します。
worksheet_name	string	書き込むワークシートの名前。
launch_application	フラグ	Excel が結果のファイルで呼び出されるかどうかを指定します。Excel を起動するパスは、「ヘルパー・アプリケーション」ダイアログ・ボックス (「ツール」メニューから「ヘルパー・アプリケーション」) 内で指定する必要があります。
generate_import	フラグ	出力されたデータ・ファイルを読み込む Excel 入力ノードが生成されるかどうかを指定します。

extensionexportnode プロパティ



拡張のエクスポート・ノードを使用すると、R スクリプトまたは Python for Spark スクリプトを実行して、データをエクスポートできます。

Python for Spark の例

```
##### script example for Python for Spark
import modeler.api
stream = modeler.script.stream()
node = stream.create("extension_export", "extension_export")
node.setPropertyValue("syntax_type", "Python")

python_script = """import spss.pyspark.runtime
from pyspark.sql import SQLContext
from pyspark.sql.types import *

cxt = spss.pyspark.runtime.getContext()
df = cxt.getSparkInputData()
print df.dtypes[:]
_newDF = df.select("Age", "Drug")
print _newDF.dtypes[:]

df.select("Age", "Drug").write.save("c:/data/ageAndDrug.json", format="json")
"""

node.setPropertyValue("python_syntax", python_script)
```

R の例

```
##### script example for R
node.setPropertyValue("syntax_type", "R")
node.setPropertyValue("r_syntax", """write.csv(modelerData, "C:/export.csv")""")
```

表 250. *extensionexportnode* プロパティ

extensionexportnode プロパティ	データ・タイプ	プロパティの説明
syntax_type	R Python	R または Python のどちらのスクリプトを実行するか指定します (R がデフォルトです)。
r_syntax	string	実行する R スクリプト・シンタックス。
python_syntax	string	実行する Python スクリプト・シンタックス。
convert_flags	StringsAndDoubles LogicalValues	フラグ型フィールドを変換するためのオプション。
convert_missing	フラグ	欠損値を R の NA 値に変換するオプションです。
convert_datetime	フラグ	日付形式または日付/時刻形式の変数を R の日付/時刻形式に変換するためのオプション。
convert_datetime_class	POSIXct POSIXlt	日付形式または日付/時刻形式の変数のうち、どの形式の変数を変換するかを指定するためのオプション。

jsonexportnode のプロパティ



JSON エクスポート・ノードは、JSON 形式でデータを出力します。

表 251. *jsonexportnode* のプロパティ

jsonexportnode のプロパティ	データ・タイプ	プロパティの説明
full_filename	string	パスを含む、完全なファイル名。
string_format	レコード 値	JSON スtring のフォーマットを指定します。デフォルトは records です。
generate_import	フラグ	エクスポートされたデータ・ファイルを読み込む JSON インポート・ノードが生成されるかどうかを指定します。デフォルトは False です。

outputfilenode プロパティ



フラット・ファイル・エクスポート・ノードでは、データが区切り文字で区切られたテキスト・ファイルへ出力されます。他の分析ソフトウェアまたはスプレッドシート・ソフトウェアで読み取ることができるデータをエクスポートする場合に役立ちます。

例

```
stream = modeler.script.stream()
outputfile = stream.createAt("outputfile", "File Output", 200, 200)
outputfile.setPropertyValue("full_filename", "c:/output/flatfile_output.txt")
outputfile.setPropertyValue("write_mode", "Append")
outputfile.setPropertyValue("inc_field_names", False)
outputfile.setPropertyValue("use_newline_after_records", False)
outputfile.setPropertyValue("delimit_mode", "Tab")
outputfile.setPropertyValue("other_delimiter", ",")
outputfile.setPropertyValue("quote_mode", "Double")
outputfile.setPropertyValue("other_quote", "*")
outputfile.setPropertyValue("decimal_symbol", "Period")
outputfile.setPropertyValue("generate_import", True)
```

表 252. *outputfilenode* プロパティ

outputfilenode プロパティ	データ・タイプ	プロパティの説明
full_filename	string	出力ファイルの名前。
write_mode	Overwrite Append	
inc_field_names	フラグ	
use_newline_after_records	フラグ	
delimit_mode	Comma Tab Space Other	
other_delimiter	CHAR	
quote_mode	None Single Double Other	
other_quote	フラグ	
generate_import	フラグ	
encoding	StreamDefault SystemDefault "UTF-8"	

sasexportnode プロパティ



SAS エクスポート・ノードで、SAS または SAS 互換ソフトウェア・パッケージで読み込むデータを、SAS 形式で出力できます。3 つの SAS ファイル・フォーマット (SAS for Windows/OS2、SAS for UNIX、または SAS バージョン 7/8) が使用可能です。

例

```
stream = modeler.script.stream()
sasexportnode = stream.createAt("sasexport", "SAS Export", 200, 200)
sasexportnode.setPropertyValue("full_filename", "c:/output/SAS_output.sas7bdat")
sasexportnode.setPropertyValue("format", "SAS8")
sasexportnode.setPropertyValue("export_names", "NamesAndLabels")
sasexportnode.setPropertyValue("generate_import", True)
```

表 253. sasexportnode プロパティ		
sasexportnode プロパティ	データ・タイプ	プロパティの説明
format	Windows	バリエーション・プロパティ・ラベル・フィールド。
	UNIX	
	SAS7	
	SAS8	
full_filename	string	
export_names	NamesAndLabels	エクスポート時にフィールド名を IBM スポス モデラー から IBM SPSS 統計 または SAS 変数名に関連付けます。
	NamesAsLabels	
generate_import	フラグ	

statisticsexportnode プロパティ



Statistics エクスポート・ノードでは、IBM SPSS 統計 .sav または .zsav フォーマットでデータを出力します。 .sav ファイルまたは .zsav ファイルは、IBM SPSS 統計 Base およびその他の製品で読み取ることができます。これは、IBM スポス モデラーのキャッシュ・ファイルに使用されるフォーマットでもあります。

このノードのプロパティについては、425 ページの『[statisticsexportnode プロパティ](#)』に記載されています。

tm1odataexport ノードのプロパティ



IBM Cognos TM1 エクスポート・ノードは、Cognos TM1 データベースで読み取ることができる形式でデータをエクスポートできます。

表 254. <i>tm1odataexport</i> ノードのプロパティ		
tm1odataexport ノードのプロパティ	データ・タイプ	プロパティの説明
<code>credential_type</code>	<code>inputCredential</code> または <code>storedCredential</code>	資格情報のタイプを示すために使用されます。
<code>input_credential</code>	リスト	<code>credential_type</code> が <code>inputCredential</code> のときは、ドメイン、ユーザー名、およびパスワードを指定します。
<code>stored_credential_name</code>	<i>string</i>	<code>credential_type</code> が <code>storedCredential</code> の場合、C & DS サーバー上の資格情報の名前を指定します。
<code>selected_cube</code>	フィールド	データのエクスポート先のキューブの名前。以下に例を示します。 <code>TM1_export.setPropertyValue("selected_cube", "plan_BudgetPlan")</code>

表 254. tm1odataexport ノードのプロパティ (続き)

tm1odataexport ノードのプロパティ	データ・タイプ	プロパティの説明
spss_field_to_tm1_element_mapping	リスト	マップされる tm1 要素は、選択されたキューブビューの列ディメンションの一部でなければなりません。形式は次のとおりです。 <pre>[[[Field_1, Dimension_1, False], [Element_1, Dimension_2, True], ...], [[Field_2, ExistMeasureElement, False], [Field_3, NewMeasureElement, True], ...]]</pre> マッピング情報を記述する 2 つのリストがあります。リーフ要素のディメンションへのマッピングは、以下の例 2 に対応します。 例 1: 最初のリスト: ([[Field_1, Dimension_1, False], [Element_1, Dimension_2, True], ...]) は、TM1 ディメンション・マップ情報に使用されます。 3 つの値を持つそれぞれのリストは、ディメンションマッピング情報を示します。3 番目のブール値は、ディメンションの要素を選択するかどうかを示すために使用されます。例えば、"[Field_1, Dimension_1, False]" は、Field_1 が Dimension_1 にマップされることを意味し、"[Element_1, Dimension_2, True]" は、Dimension_2 に対して Element_1 が選択されることを意味します。 例 2: 2 番目のリスト: ([[Field_2, ExistMeasureElement, False], [Field_3, NewMeasureElement, True], ...]) は、TM1 数値データ・ディメンション要素のマップ情報に使用されます。 3 つの値を持つそれぞれのリストは、測定要素のマッピング情報を示します。3 番目のブール値は、新しい要素を作成する必要があることを示すために使用されます。"[Field_2, ExistMeasureElement, False]" は、Field_2 が ExistMeasureElement にマップされることを示します。"[Field_3, NewMeasureElement, True]" は、NewMeasureElement が selected_measure で選択された数値データディメンションである必要があり、Field_3 がそれにマップされることを示します。
selected_measure	string	数値データ・ディメンションを指定します。 例: <pre>setProperty("selected_measure", "Measures")</pre>
connection_type	AdminServer TM1Server	接続タイプを指定します。デフォルトは AdminServer です。

表 254. tm1odataexport ノードのプロパティ (続き)

tm1odataexport ノードのプロパティ	データ・タイプ	プロパティの説明
admin_host	string	REST API のホスト名の URL。 connection_type が AdminServer の場合は必須です。
server_name	string	admin_host から選択した TM1 サーバーの名前。 connection_type が AdminServer の場合は必須です。
server_url	string	TM1 サーバー REST API の URL。 connection_type が TM1Server の場合は必須です。

tm1export ノードのプロパティ (廃止)



IBM Cognos TM1 エクスポート・ノードは、Cognos TM1 データベースで読み取ることができる形式でデータをエクスポートできます。

注: このノードは、Modeler 18.0 で廃止されました。それに置き換わるノードのスクリプト名は tm1odataexport です。

表 255. tm1export ノードのプロパティ

tm1export ノードのプロパティ	データ・タイプ	プロパティの説明
pm_host	string	注: バージョン 16.0 および 17.0 の場合のみ ホスト名。例: TM1_export.setPropertyValue("pm_host", 'http://9.191.86.82:9510/pmhub/pm')
tm1_connection	["フィールド", "フィールド", ..., "フィールド"]	注: バージョン 16.0 および 17.0 の場合のみ TM1 サーバーの接続の詳細を含むリストのプロパティ。形式は次のとおりです。 ["TM1_Server_Name", "tm1_username", "tm1_password"] 例: TM1_export.setPropertyValue("tm1_connection", ['Planning Sample', "admin" "apple"])
selected_cube	フィールド	データのエクスポート先のキューブの名前。 例: TM1_export.setPropertyValue("selected_cube", "plan_BudgetPlan")

表 255. tm1export ノードのプロパティ (続き)

tm1export ノードのプロパティ	データ・タイプ	プロパティの説明
spssfield_tm1element_mapping	リスト	マップされる tm1 要素は、選択されたキューブビューの列ディメンションの一部でなければなりません。形式は次のとおりです。 <pre>[[[Field_1, Dimension_1, False], [Element_1, Dimension_2, True], ...], [[Field_2, ExistMeasureElement, False], [Field_3, NewMeasureElement, True], ...]]</pre> マッピング情報を示す 2 つのリストがあります。ディメンションへの葉要素のマッピングは、以下の例 2 に対応しています。 例 1: 最初のリスト: ([[Field_1, Dimension_1, False], [Element_1, Dimension_2, True], ...]) は、TM1 ディメンションのマッピング情報に使用されます。 3 つの値を持つそれぞれのリストは、ディメンションマッピング情報を示します。3 番目のブール値は、ディメンションの要素を選択するかどうかを示すために使用されます。例えば、"[Field_1, Dimension_1, False]" は Field_1 が Dimension_1 にマップされることを意味し、"[Element_1, Dimension_2, True]" は Element_1 が Dimension_2 に選択されることを意味します。 例 2: 2 番目のリスト ([[Field_2, ExistMeasureElement, False], [Field_3, NewMeasureElement, True], ...]) は、TM1 数値データ・ディメンション要素のマッピング情報に使用されます。 3 つの値を持つそれぞれのリストは、測定要素のマッピング情報を示します。3 番目のブール値は、新しい要素を作成する必要があることを示すために使用されます。"[Field_2, ExistMeasureElement, False]" は、Field_2 が ExistMeasureElement にマップされることを意味します。"[Field_3, NewMeasureElement, True]" は、NewMeasureElement が selected_measure で選択された数値データ・ディメンションである必要があり、Field_3 がそれにマップされることを意味します。

表 255. *tm1export* ノードのプロパティ (続き)

tm1export ノードのプロパティ	データ・タイプ	プロパティの説明
selected_measure	string	数値データ ディメンションを指定します。 例:setPropertyValue("selected_measure", "Measures")

xmlexportnode プロパティ



XML エクスポート・ノードでは、XML 形式のファイルにデータを出力します。オプションで、エクスポートしたデータをストリームに読み込む XML 入力ノードを作成できます。

例

```
stream = modeler.script.stream()
xmlexportnode = stream.createAt("xmlexport", "XML Export", 200, 200)
xmlexportnode.setPropertyValue("full_filename", "c:/export/data.xml")
xmlexportnode.setPropertyValue("map", ["/catalog/book/genre", "genre"], ["/catalog/book/title", "title"])
```

表 256. *xmlexportnode* プロパティ

xmlexportnode プロパティ	データ・タイプ	プロパティの説明
full_filename	string	(必須) XML エクスポート・ファイルの完全パスおよびファイル名。
use_xml_schema	フラグ	XML スキーマ (XSD ファイルまたは DTD ファイル) を使用して、エクスポートされたデータの構造を制御するかどうかを指定します。
full_schema_filename	string	使用する XSD ファイルまたは DTD ファイルの完全パスおよびファイル名。 <i>use_xml_schema</i> が true に設定されている場合にのみ必須です。
generate_import	フラグ	エクスポートされたデータ・ファイルをストリームに読み込む XML 入力ノードを、自動的に生成します。
records	string	レコードの境界を示す XPath 式。
map	string	XML 構造にフィールド名をマッピングします。

第 18 章 IBM SPSS 統計 ノードのプロパティ

statisticsimportnode プロパティ



Statistics ファイル・ノードは、同じ形式を使用する IBM SPSS 統計 で使用される .sav または .zsav ファイル形式のデータおよび IBM スポス モデラー に保存されたキャッシュ・ファイルを読み込みます。

例

```
stream = modeler.script.stream()
statisticsimportnode = stream.createAt("statisticsimport", "SAV Import",
200, 200)
statisticsimportnode.setPropertyValue("full_filename", "C:/data/drug1n.sav")
statisticsimportnode.setPropertyValue("import_names", True)
statisticsimportnode.setPropertyValue("import_data", True)
```

表 257. statisticsimportnode プロパティ

statisticsimportnode プロパティ	データ・タイプ	プロパティの説明
full_filename	string	パスを含む、完全なファイル名。
password	string	パスワード。 password パラメータは、file_encrypted パラメータよりも前に設定する必要があります。
file_encrypted	フラグ	ファイルがパスワード保護されているかどうか。
import_names	NamesAndLabels LabelsAsNames	変数名と変数ラベルを処理する方法。
import_data	DataAndLabels LabelsAsData	値とラベルを処理する方法。
use_field_format_for_storage	Boolean	インポート時に IBM SPSS 統計 フィールド形式情報を使用するかどうかを指定します。

statistictransformnode プロパティ



Statistics 変換ノードは、IBM スポス モデラー のデータ・ソースに対する IBM SPSS 統計 シンタックス・コマンドの選択を行います。 このノードには、ライセンス交付を受けた IBM SPSS 統計 のコピーが必要です。

例

```
stream = modeler.script.stream()
statistictransformnode = stream.createAt("statistictransform",
"Transform", 200, 200)
statistictransformnode.setPropertyValue("syntax", "COMPUTE NewVar = Na +
```

```
K.")
statisticstransformnode.setKeyedPropertyValue("new_name", "NewVar", "Mixed
Drugs")
statisticstransformnode.setPropertyValue("check_before_saving", True)
```

表 258. *statisticstransformnode* プロパティ

statisticstransformnode プロパティ	データ・タイプ	プロパティの説明
syntax	string	
check_before_saving	フラグ	項目を保存する前に、入力されたシンタックスを検証します。シンタックスが無効な場合は、エラー・メッセージを表示します。
default_include	フラグ	詳しくは、トピック 165 ページの『 filternode プロパティ 』を参照してください。
include	フラグ	詳しくは、トピック 165 ページの『 filternode プロパティ 』を参照してください。
new_name	string	詳しくは、トピック 165 ページの『 filternode プロパティ 』を参照してください。

statisticsmodelnode プロパティ



Statistics モデル・ノードを使用すると、PMML を作成する IBM SPSS 統計手続きを実行してデータを分析および使用することができます。このノードには、ライセンス交付を受けた IBM SPSS 統計のコピーが必要です。

例

```
stream = modeler.script.stream()
statisticsmodelnode = stream.createAt("statisticsmodel", "Model", 200, 200)
statisticsmodelnode.setPropertyValue("syntax", "COMPUTE NewVar = Na + K.")
statisticsmodelnode.setKeyedPropertyValue("new_name", "NewVar", "Mixed
Drugs")
```

statisticsmodelnode プロパティ	データ・タイプ	プロパティの説明
syntax	string	
default_include	フラグ	詳しくは、トピック 165 ページの『 filternode プロパティ 』を参照してください。
include	フラグ	詳しくは、トピック 165 ページの『 filternode プロパティ 』を参照してください。
new_name	string	詳しくは、トピック 165 ページの『 filternode プロパティ 』を参照してください。

statisticsoutputnode プロパティー



Statistics 出力ノードを使用すると、IBM SPSS 統計 手続きを呼び出し、IBM スポス モデラー データを分析することができます。さまざまな IBM SPSS 統計 分析手続き にアクセスできます。このノードには、ライセンス交付を受けた IBM SPSS 統計のコピーが必要です。

例

```
stream = modeler.script.stream()
statisticsoutputnode = stream.createAt("statisticsoutput", "Output", 200,
200)
statisticsoutputnode.setPropertyValue("syntax", "SORT CASES BY Age(A) Sex(A)
BP(A) Cholesterol(A)")
statisticsoutputnode.setPropertyValue("use_output_name", False)
statisticsoutputnode.setPropertyValue("output_mode", "File")
statisticsoutputnode.setPropertyValue("full_filename", "Cases by Age, Sex
and Medical History")
statisticsoutputnode.setPropertyValue("file_type", "HTML")
```

表 259. statisticsoutputnode プロパティー

statisticsoutputnode プロパティー	データ・タイプ	プロパティーの説明
mode	Dialog Syntax	「IBM SPSS 統計 ダイアログ」オ プションまたはシンタックス・ エディターを選択します。
syntax	string	
use_output_name	フラグ	
output_name	string	
output_mode	Screen File	
full_filename	string	
file_type	HTML SPV SPW	

statisticsexportnode プロパティー



Statistics エクスポート・ノードでは、IBM SPSS 統計 .sav または .zsav フォーマット
でデータを出力します。 .sav ファイルまたは .zsav ファイルは、 IBM SPSS 統計
Base およびその他の製品で読み取ることができます。これは、IBM スポス モデラー
のキャッシュ・ファイルに使用されるフォーマットでもあります。

例

```
stream = modeler.script.stream()
statisticsexportnode = stream.createAt("statisticsexport", "Export", 200,
200)
statisticsexportnode.setPropertyValue("full_filename", "c:/output/
```

```

SPSS_Statistics_out.sav")
statisticsexportnode.setPropertyValue("field_names", "Names")
statisticsexportnode.setPropertyValue("launch_application", True)
statisticsexportnode.setPropertyValue("generate_import", True)

```

表 260. *statisticsexportnode* プロパティ

statisticsexportnode プロパティ	データ・タイプ	プロパティの説明
full_filename	<i>string</i>	
file_type	sav zsav	ファイルを <i>sav</i> または <i>zsav</i> 形式で保存します。以下に例を示します。 <code>statisticsexportnode.setPropertyValue("file_type", "sav")</code>
encrypt_file	フラグ	ファイルがパスワード保護されているかどうか。
password	<i>string</i>	パスワード。
launch_application	フラグ	
export_names	NamesAndLabels NamesAsLabels	エクスポート時にフィールド名を IBM スポス モデラー から IBM SPSS 統計 または SAS 変数名に関連付けます。
generate_import	フラグ	

第 19 章 Python ノードのプロパティ

gmm のプロパティ



ガウス混合[®]モデルは、未知パラメータを持つ有限個数のガウス分布の混合からすべてのデータ ポイントが生成されると仮定する確率モデルです。混合モデルは、データの共分散構造および潜在ガウス分布の中心に関する情報を取り込むための一般化 K-means クラスタリングと考えることができます。スポンサー モデラー のガウス混合 ノードは、ガウス混合ライブラリーのコア機能およびよく使用されるパラメータを公開します。このノードは Python で実装されています。

表 261. gmm のプロパティ

gmm のプロパティ	データ・タイプ	プロパティの説明
use_partition	<i>Boolean</i>	True または False に設定して、データ区分データを使用するかどうかを指定します。デフォルトは False です。
covariance_type	<i>string</i>	Full、Tied、Diag、または Spherical を指定して共分散タイプを設定します。
number_component	<i>integer</i>	混合コンポーネントの数の整数を指定します。最小値は 1 です。デフォルト値は 2 です。
component_lable	<i>Boolean</i>	クラスター ラベルを文字列に設定するには True を指定し、クラスター ラベルを数値に設定するには False を指定します。デフォルトは False です。
label_prefix	<i>string</i>	文字列のクラスター ラベルを使用する場合は、接頭辞を指定できます。
enable_random_seed	<i>Boolean</i>	ランダム シードを使用にする場合は、True を指定します。デフォルトは False です。
random_seed	<i>integer</i>	ランダム シードを使用する場合は、無作為サンプルの生成に使用する整数を指定します。
tol	<i>DOUBLE</i>	収束のしきい値を指定します。デフォルトは 0.000.1 です。
max_iter	<i>integer</i>	実行する反復の最大回数を指定します。デフォルトは 100 です。
init_params	<i>string</i>	使用する初期設定パラメーターを設定します。オプションは Kmeans または Random です。
warm_start	<i>Boolean</i>	最後の適合の解を、適合の次の呼び出しの初期設定として使用するには、True を指定します。デフォルトは False です。

hdbscannode のプロパティ



Hierarchical Density-Based Spatial Clustering (HDBSCAN)[®] は、教師なし学習を使用してデータ・セットのクラスター (つまり、密度の高い領域) を検出します。スボスモデラーの HDBSCAN ノードは、HDBSCAN ライブラリーのコア機能およびよく使用されるパラメーターを公開します。このノードは Python で実装されており、最初にグループの性質が分からない場合にデータ・セットを異なるグループにクラスター化するために使用できます。

表 262. hdbscannode のプロパティ

hdbscannode のプロパティ	データ・タイプ	プロパティの説明
inputs	フィールド	クラスタリングの入力フィールド。
useHPO	Boolean	Rbfopt に基づくハイパーパラメータ最適化を有効にするには true を指定し、無効にするには false を指定します。ハイパーパラメータ最適化を有効にすると、パラメータの最適な組み合わせが自動的に検出され、サンプルに対するモデルの誤差率が予測値以下になります。デフォルトは false です。
min_cluster_size	integer	クラスターの最小サイズ。整数を指定してください。デフォルトは 5 です。
min_samples	integer	あるポイントがコアポイントと見なされるための近傍のサンプル数。整数を指定してください。0 に設定した場合、min_cluster_size が使用されます。デフォルトは 0 です。
algorithm	string	使用するアルゴリズムを best、generic、prims_kdtree、prims_balltree、boruvka_kdtree、または boruvka_balltree から指定します。デフォルトは best です。
metric	string	機能配列内のインスタンス間の距離を計算するときに使用するメトリックを euclidean、cityblock、L1、L2、manhattan、braycurtis、canberra、chebyshev、correlation、minkowski、または sqeuclidean から指定します。デフォルトは euclidean です。
useStringLabel	Boolean	文字列のクラスターラベルを使用するには true を指定し、数値のクラスターラベルを使用するには false を指定します。デフォルトは false です。
stringLabelPrefix	string	useStringLabel パラメーターが true に設定されている場合は、文字列のラベル接頭辞の値を指定します。デフォルトの接頭辞は cluster です。

表 262. *hdbscannode* のプロパティ (続き)

hdbscannode のプロパティ	データ・タイプ	プロパティの説明
<code>approx_min_span_tree</code>	<i>Boolean</i>	近似最小スパンニングツリーを受け入れるには <code>true</code> を指定します。正確さのために速度を犠牲にする場合は <code>false</code> を指定します。デフォルトは <code>true</code> です。
<code>cluster_selection_method</code>	<i>string</i>	圧縮ツリーからクラスターを選択するために使用するメソッドを <code>eom</code> または <code>leaf</code> から指定します。デフォルトは <code>eom</code> (Excess of Mass アルゴリズム) です。
<code>allow_single_cluster</code>	<i>Boolean</i>	単一クラスター結果を許可する場合は、 <code>true</code> を指定します。デフォルトは <code>false</code> です。
<code>p_value</code>	<i>DOUBLE</i>	メトリックに <code>minkowski</code> を使用している場合は、使用する <code>p_value</code> を指定します。デフォルトは <code>1.5</code> です。
<code>leaf_size</code>	<i>integer</i>	スペース ツリー アルゴリズム (<code>boruvka_kdtree</code> または <code>boruvka_balltree</code>) を使用している場合は、ツリーのリーフ ノード内のポイント数を指定します。デフォルトは <code>40</code> です。
<code>outputValidity</code>	<i>Boolean</i>	妥当性インデックス グラフがモデル出力に含まれるかどうかを制御するには <code>true</code> または <code>false</code> を指定します。
<code>outputCondensed</code>	<i>Boolean</i>	圧縮ツリー グラフがモデル出力に含まれるかどうかを制御するには <code>true</code> または <code>false</code> を指定します。
<code>outputSingleLinkage</code>	<i>Boolean</i>	単結合ツリー グラフがモデル出力に含まれるかどうかを制御するには <code>true</code> または <code>false</code> を指定します。
<code>outputMinSpan</code>	<i>Boolean</i>	最小スパンニング ツリー グラフがモデル出力に含まれるかどうかを制御するには <code>true</code> または <code>false</code> を指定します。
<code>is_split</code>		バージョン 18.2.1.1 で追加されました。

kdemodel のプロパティ



カーネル密度推定 (KDE)[®] は、Ball Tree または KD Tree のアルゴリズムを使用してクエリを効率化し、教師なし学習、特徴量エンジニアリング、データのモデル化の概念を結合します。KDE などの近隣ベースの手法が、最もよく使用され、有用な密度推定手法です。スボス モデラー の KDE モデル作成および KDE シミュレーションのノードは、KDE ライブラリーのコア機能およびよく使用されるパラメータを公開します。これらのノードは Python で実装されています。

表 263. *kdemodel* のプロパティ

kdemodel のプロパティ	データ・タイプ	プロパティの説明
bandwidth	<i>DOUBLE</i>	デフォルトは 1 です。
kernel	<i>string</i>	使用するカーネル: gaussian、tophat、epanechnikov、exponential、linear、または cosine。デフォルトは gaussian です。
algorithm	<i>string</i>	使用するツリー・アルゴリズム: kd_tree、ball_tree、または auto。デフォルトは auto です。
metric	<i>string</i>	距離を計算するときに使用するメトリック。kd_tree アルゴリズムの場合は、Euclidean、Chebyshev、Cityblock、Minkowski、Manhattan、Infinity、P、L2、または L1 から選択します。ball_tree アルゴリズムの場合は、Euclidian、Braycurtis、Chebyshev、Canberra、Cityblock、Dice、Hamming、Infinity、Jaccard、L1、L2、Minkowski、Matching、Manhattan、P、Rogersanimoto、Russellrao、Sokalmichener、Sokalsneath、または Kulsinski から選択します。デフォルトは Euclidean です。
atol	<i>Float</i>	希望する結果の絶対許容度。一般に、許容度を大きくすると、実行速度が上がります。デフォルトは 0.0 です。
rtol	<i>Float</i>	希望する結果の相対許容度。一般に、許容度を大きくすると、実行速度が上がります。デフォルトは 1E-8 です。
breadthFirst バージョン 18.2.1.1 以降は breadth_first に名前変更されています。	<i>Boolean</i>	幅優先のアプローチを使用するには、True に設定します。深さ優先のアプローチを使用するには、False に設定します。デフォルトは True です。
LeafSize バージョン 18.2.1.1 以降は leaf_size に名前変更されています。	<i>integer</i>	基本となるツリーのリーフサイズ。デフォルトは 40 です。この値を変更すると、パフォーマンスに大きな影響を与える可能性があります。
pValue	<i>DOUBLE</i>	メトリックに Minkowski を使用している場合は、使用する P 値を指定します。デフォルトは 1.5 です。
custom_name		
default_node_name		

表 263. *kdemodel* のプロパティ (続き)

kdemodel のプロパティ	データ・タイプ	プロパティの説明
use_HPO		

kdeexport プロパティ



カーネル密度推定 (KDE)[®] は、Ball Tree または KD Tree のアルゴリズムを使用してクエリを効率化し、教師なし学習、特徴量エンジニアリング、データのモデル化の概念を結合します。KDE などの近隣ベースの手法が、最もよく使用され、有用な密度推定手法です。スボス モデラー の KDE モデル作成および KDE シミュレーションのノードは、KDE ライブラリーのコア機能およびよく使用されるパラメータを公開します。これらのノードは Python で実装されています。

表 264. *kdeexport* プロパティ

kdeexport プロパティ	データ・タイプ	プロパティの説明
bandwidth	<i>DOUBLE</i>	デフォルトは 1 です。
kernel	<i>string</i>	使用するカーネル: gaussian または tophat。デフォルトは gaussian です。
algorithm	<i>string</i>	使用するツリー・アルゴリズム: kd_tree、ball_tree、または auto。デフォルトは auto です。
metric	<i>string</i>	距離を計算するときに使用するメトリック。kd_tree アルゴリズムの場合は、Euclidean、Chebyshev、Cityblock、Minkowski、Manhattan、Infinity、P、L2、または L1 から選択します。ball_tree アルゴリズムの場合は、Euclidian、Braycurtis、Chebyshev、Canberra、Cityblock、Dice、Hamming、Infinity、Jaccard、L1、L2、Minkowski、Matching、Manhattan、P、Rogersanimoto、Russellrao、Sokalmichener、Sokalsneath、または Kulsinski から選択します。デフォルトは Euclidean です。
atol	<i>Float</i>	希望する結果の絶対許容度。一般に、許容度を大きくすると、実行速度が上がります。デフォルトは 0.0 です。
rtol	<i>Float</i>	希望する結果の相対許容度。一般に、許容度を大きくすると、実行速度が上がります。デフォルトは 1E-8 です。
breadthFirst	<i>Boolean</i>	幅優先のアプローチを使用するには、True に設定します。深さ優先のアプローチを使用するには、False に設定します。デフォルトは True です。

表 264. kdeexport プロパティ (続き)

kdeexport プロパティ	データ・タイプ	プロパティの説明
LeafSize	integer	基本となるツリーのリーフ サイズ。デフォルトは 40 です。この値を変更すると、パフォーマンスに大きな影響を与える可能性があります。
pValue	DOUBLE	メトリックに Minkowski を使用している場合は、使用する P 値を指定します。デフォルトは 1.5 です。

gmm のプロパティ



ガウス混合[®]モデルは、未知パラメータを持つ有限個数のガウス分布の混合からすべてのデータ ポイントが生成されると仮定する確率モデルです。混合モデルは、データの共分散構造および潜在ガウス分布の中心に関する情報を取り込むための一般化 K-means クラスタリングと考えることができます。スポンサー モデラー のガウス混合 ノードは、ガウス混合ライブラリーのコア機能およびよく使用されるパラメータを公開します。このノードは Python で実装されています。

表 265. gmm のプロパティ

gmm のプロパティ	データ・タイプ	プロパティの説明
use_partition	Boolean	True または False に設定して、データ区分データを使用するかどうかを指定します。デフォルトは False です。
covariance_type	string	Full、Tied、Diag、または Spherical を指定して共分散タイプを設定します。
number_component	integer	混合コンポーネントの数の整数を指定します。最小値は 1 です。デフォルト値は 2 です。
component_lable	Boolean	クラスター ラベルを文字列に設定するには True を指定し、クラスター ラベルを数値に設定するには False を指定します。デフォルトは False です。
label_prefix	string	文字列のクラスター ラベルを使用する場合は、接頭辞を指定できます。
enable_random_seed	Boolean	ランダム シードを使用にする場合は、True を指定します。デフォルトは False です。
random_seed	integer	ランダム シードを使用する場合は、無作為サンプルの生成に使用する整数を指定します。
tol	DOUBLE	収束のしきい値を指定します。デフォルトは 0.000.1 です。
max_iter	integer	実行する反復の最大回数を指定します。デフォルトは 100 です。

表 265. gmm のプロパティ (続き)

gmm のプロパティ	データ・タイプ	プロパティの説明
init_params	string	使用する初期設定パラメーターを設定します。オプションは Kmeans または Random です。
warm_start	Boolean	最後の適合の解を、適合の次の呼び出しの初期設定として使用するには、True を指定します。デフォルトは False です。

ocsvmnode のプロパティ



One-Class SVM ノードでは、教師なし学習アルゴリズムを使用します。このノードは、新規性検知の目的で使用できます。このノードは、与えられたサンプル・セットのソフト境界を検知し、新規ポイントがこのセットに属するか、属さないかを分類します。スボス モデラー の One-Class SVM モデル作成ノードは Python に実装されており、scikit-learn® Python ライブラリーを必要とします。

表 266. ocsvmnode のプロパティ

ocsvmnode のプロパティ	データ・タイプ	プロパティの説明
role_use バージョン 18.2.1.1 以降は custom_fields に名前変更されています。	string	定義済みの役割を使用するには predefined を指定し、ユーザー設定フィールドの割り当てを使用するには custom を指定します。デフォルトは predefined です。
splits	フィールド	分割用のフィールド名のリスト。
use_partition	Boolean	true または false を指定します。デフォルトは true です。true に設定した場合は、モデルの構築時に学習データのみが使用されます。
mode_type	string	モード。指定できる値は、simple または expert です。simple を指定した場合は、「エキスパート」タブのすべてのパラメーターが無効になります。
stopping_criteria	string	指数表記の文字列。指定できる値は、1.0E-1、1.0E-2、1.0E-3、1.0E-4、1.0E-5、または 1.0E-6 です。デフォルトは 1.0E-3 です。
precision	Float	回帰精度 (ニュー)。学習誤差およびサポート・ベクターの小数部の範囲です。0 より大きく 1.0 以下の数値を指定してください。デフォルトは 0.1 です。
kernel	string	アルゴリズムで使用するカーネルタイプ。指定できる値は linear、poly、rbf、sigmoid、または precomputed です。デフォルトは rbf です。

表 266. ocsvmnode のプロパティ (続き)

ocsvmnode のプロパティ	データ・タイプ	プロパティの説明
enable_gamma	Boolean	gamma パラメータを有効にします。true または false を指定します。デフォルトは true です。
gamma	Float	このパラメータはカーネル rbf、poly、および sigmoid の場合にのみ有効です。enable_gamma パラメータを false に設定した場合、このパラメータは auto に設定されます。true に設定した場合、デフォルトは 0.1 です。
coef0	Float	カーネル関数の独立した項目。このパラメータは、poly カーネルおよび sigmoid カーネルの場合にのみ有効です。デフォルト値は 0.0 です。
degree	integer	多項式のカーネル関数の次数。このパラメータは、poly カーネルの場合にのみ有効です。任意の整数を指定します。デフォルトは 3 です。
shrinking	Boolean	収縮ヒューリスティック・オプションを使用するかどうかを指定します。true または false を指定します。デフォルトは false です。
enable_cache_size	Boolean	cache_size パラメータを有効にします。true または false を指定します。デフォルトは false です。
cache_size	Float	カーネル キャッシュのサイズ (MB)。デフォルトは 200 です。
pc_type	string	平行座標グラフィックスのタイプ。指定できるオプションは independent または general です。
lines_amount	integer	グラフィックに含める最大行数。1 から 1000 の整数を指定します。
lines_fields_custom	Boolean	lines_fields パラメータを有効にし、グラフ出力にカスタム・フィールドを表示できるようにします。false に設定した場合、すべてのフィールドが表示されます。true に設定した場合、lines_fields パラメータで指定したフィールドのみが表示されます。パフォーマンス上の理由から、最大で 20 個のフィールドが表示されます。
lines_fields	フィールド	グラフィックに垂直軸として含めるフィールド名のリスト。

表 266. ocsvmnode のプロパティ (続き)

ocsvmnode のプロパティ	データ・タイプ	プロパティの説明
enable_graphic	Boolean	true または false を指定します。グラフィック出力を有効にします (時間を節約し、ストリーム・ファイル・サイズを削減したい場合は、このオプションを無効にしてください)。
enable_hpo	Boolean	true または false に指定して、HPO オプションを有効または無効にします。true に設定すると Rbfopt が適用され、自動的に「最適な」One-Class SVM モデルが判別されます。この場合、ユーザーが以下の target_objval パラメーターで定義した目標値に到達します。
target_objval	Float	目標とする目的関数 (サンプルに対するモデルの誤差率) の値 (例えば、未知の最適条件の値)。最適条件が不明な場合は、このパラメーターを適当な値 (0.01 など) に設定してください。
max_iterations	integer	モデルを試行する最大反復数。デフォルトは 1000 です。
max_evaluations	integer	速度より精度を重視する場合の、モデルを試行するための関数評価の最大回数。デフォルトは 300 です。

rfnode プロパティ



ランダム・フォレスト・ノードは、ツリー・モデルを基本モデルとして使用するバギング・アルゴリズムの高度な実装を使用します。スポンサー モデラー のランダムフォレストモデル作成ノードは Python に実装されており、scikit-learn® Python ライブラリーを必要とします。

表 267. rfnode プロパティ

rfnode プロパティ	データ・タイプ	プロパティの説明
role_use	string	定義済みの役割を使用するには predefined を指定し、ユーザー設定フィールドの割り当てを使用するには custom を指定します。デフォルトは predefined です。
inputs	フィールド	入力用のフィールド名のリスト。
splits	フィールド	分割用のフィールド名のリスト。
n_estimators	integer	作成するツリーの数。デフォルトは 10 です。

表 267. rfnode プロパティ (続き)

rfnode プロパティ	データ・タイプ	プロパティの説明
specify_max_depth	Boolean	カスタムの最大の深さを指定します。 false の場合、リーフがすべて純粋なリーフになるまで、またはすべてのリーフのサンプル数が min_samples_split 未満になるまでノードが展開されます。デフォルトは false です。
max_depth	integer	ツリーの最大の深さ。デフォルトは 10 です。
min_samples_leaf	integer	リーフ ノードの最小サイズ。デフォルトは 1 です。
max_features	string	最良の分割を求めるときに考慮するフィーチャーの数。 • auto の場合、分類に max_features=sqrt(n_features) を使用し、回帰に max_features=sqrt(n_features) を使用します。 • sqrt の場合は max_features=sqrt(n_features) です。 • log2 の場合は max_features=log2(n_features) です。 デフォルトは auto です。
bootstrap	Boolean	ツリーの作成時にブートストラップサンプルを使用します。デフォルトは true です。
oob_score	Boolean	Out of Bag サンプルを使用して一般化の精度を推定します。デフォルト値は false です。
extreme	Boolean	Extremely Randomized Trees を使用します。デフォルトは false です。
use_random_seed	Boolean	結果を再現するには、これを指定します。デフォルトは false です。
random_seed	integer	ツリーの作成時に使用する乱数シード。任意の整数を指定します。
cache_size	Float	カーネル キャッシュのサイズ (MB)。デフォルトは 200 です。
enable_random_seed	Boolean	random_seed パラメータを有効にします。true または false を指定します。デフォルトは false です。

表 267. rfnode プロパティ (続き)

rfnode プロパティ	データ・タイプ	プロパティの説明
enable_hpo	Boolean	true または false に指定して、HPO オプションを有効または無効にします。true に設定すると Rbfopt が適用され、自動的に「最適な」ランダム フォレスト モデルが判別されます。この場合、ユーザーが以下の target_objval パラメータで定義した目標値に到達します。
target_objval	Float	目標とする目的関数 (サンプルに対するモデルの誤差率) の値 (例えば、未知の最適条件の値)。最適条件が不明な場合は、このパラメータを適当な値 (0.01 など) に設定してください。
max_iterations	integer	モデルを試行する最大反復数。デフォルトは 1000 です。
max_evaluations	integer	速度より精度を重視する場合の、モデルを試行するための関数評価の最大回数。デフォルトは 300 です。

smotenode のプロパティ



SMOTE (Synthetic Minority Over-sampling Technique) ノードは不均衡データ・セットを扱うためのオーバーサンプリング・アルゴリズムを提供します。これにより、データの均衡化のための高度な手法が提供されます。スボス モデラー の SMOTE プロセス・ノードは Python に実装されており、imbalanced-learn® Python ライブラリーを必要とします。

表 268. smotenode のプロパティ

smotenode のプロパティ	データ・タイプ	プロパティの説明
target_field バージョン 18.2.1.1 以降は target に名前変更されています。	フィールド	対象フィールド。
sample_ratio	string	カスタムの比率の値を使用できるようにします。オプションは、自動 (sample_ratio_auto) と比率の設定 (sample_ratio_manual) の 2 つです。
sample_ratio_value	Float	これは、マジョリティー クラスのサンプル数に対するマイノリティー クラスのサンプル数の比率です。0 より大きく 1 以下でなければなりません。デフォルトは auto です。
enable_random_seed	Boolean	true に設定した場合は random_seed プロパティが有効になります。
random_seed	integer	乱数発生ルーチンによって使用されるシード。

表 268. smotenode のプロパティ (続き)

smotenode のプロパティ	データ・タイプ	プロパティの説明
k_neighbours	integer	合成サンプルを作成するために使用する最近傍の数。デフォルトは 5 です。
m_neighbours	integer	マイノリティー サンプルが危険な状況にあるかをどうか判断するために使用する最近傍の数。このオプションは、SMOTE アルゴリズムの種類が borderline1 および borderline2 の場合にのみ使用可能です。デフォルトは 10 です。
algorithm_kind バージョン 18.2.1.1 以降は algorithm に名前変更されています。	string	SMOTE アルゴリズムの種類: regular 、 borderline1 、または borderline2 。
usepartition バージョン 18.2.1.1 以降は use_partition に名前変更されています。	Boolean	true に設定した場合は、モデルの構築に学習データのみが使用されます。デフォルトは true です。

tsnnode プロパティ



t 分布 Stochastic Neighbor Embedding (t-SNE) は、高次元データの視覚化のためのツールです。t-SNE は、データ ポイントの類似性を確率に変換します。スポン モデラーの t-SNE ノードは Python で実装されており、scikit-learn® Python ライブラリーを必要とします。

表 269. tsnnode プロパティ

tsnnode プロパティ	データ・タイプ	プロパティの説明
mode_type	string	simple モードまたは expert モードを指定します。
n_components	string	埋め込み空間の次元 (2 次元 または 3 次元)。2 または 3 を指定します。デフォルトは 2 です。
method	string	barnes_hut または exact を指定します。デフォルトは barnes_hut です。
init	string	埋め込みの初期化。 random または pca を指定します。デフォルトは random です。
target_field バージョン 18.2.1.1 以降は target に名前変更されています。	string	対象フィールド名。これが出力グラフのカラー マップになります。対象フィールドを指定しないと、グラフには 1 色が使用されます。

表 269. tsnenode プロパティ (続き)

tsnenode プロパティ	データ・タイプ	プロパティの説明
perplexity	Float	Perplexity は、他の多様体学習アルゴリズムで使用される最近隣の数に関連します。通常、データセットが大きいほど、必要とされる Perplexity も大きくなります。5 から 50 の間の値を選択することを考慮してください。デフォルトは 30 です。
early_exaggeration	Float	埋め込み空間における、元の空間の自然クラスタの気密度、およびクラスタ間の間隔を制御します。デフォルトは 12.0 です。
learning_rate	Float	デフォルトは 200 です。
n_iter	integer	最適化の最大反復数。250 以上に設定してください。デフォルトは 1000 です。
angle	Float	あるポイントから測定した遠方ノードの角サイズ。0 から 1 の間の値を指定します。デフォルトは 0.5 です。
enable_random_seed	Boolean	random_seed パラメータを有効にするには、true に設定します。デフォルトは false です。
random_seed	integer	使用する乱数シード。デフォルトは None です。
n_iter_without_progress	integer	進捗のない最大反復数。デフォルトは 300 です。
min_grad_norm	string	勾配ノルムがこのしきい値を下回る場合、最適化は中断されます。デフォルトは 1.0E-7 です。指定できる値は次のとおりです： • 1.0E-1 • 1.0E-2 • 1.0E-3 • 1.0E-4 • 1.0E-5 • 1.0E-6 • 1.0E-7 • 1.0E-8
isGridSearch	Boolean	複数の異なる Perplexity で t-SNE を実行するには、true に設定します。デフォルトは false です。
output_Rename	Boolean	カスタム名を設定する場合は true を指定し、出力に名前を自動的に付ける場合は false を指定します。デフォルトは false です。

表 269. tsnode プロパティ (続き)

tsnode プロパティ	データ・タイプ	プロパティの説明
output_to	string	Screen または Output を指定します。デフォルトは Screen です。
full_filename	string	出力ファイル名を指定します。
output_file_type	string	出力ファイル形式。HTML または Output object を指定します。デフォルトは HTML です。

xgboostlinearnode のプロパティ



XGBoost Linear[®] は、線型モデルを基本モデルとして使用する勾配ブースティング・アルゴリズムの高度な実装です。ブースティング・アルゴリズムでは、弱い分類子に繰り返し学習させ、それを最終的な強い分類子に追加します。スポンサー モデラー の XGBoost Linear ノードは Python で実装されています。

表 270. xgboostlinearnode のプロパティ

xgboostlinearnode のプロパティ	データ・タイプ	プロパティの説明
TargetField バージョン 18.2.1.1 以降は target に名前変更されています。	フィールド	
InputFields バージョン 18.2.1.1 以降は inputs に名前変更されています。	フィールド	
alpha	DOUBLE	アルファ線型ブースティングパラメータ。0 以上の任意の数値を指定してください。デフォルトは 0 です。
lambda	DOUBLE	ラムダ線型ブースティングパラメータ。0 以上の任意の数値を指定してください。デフォルトは 1 です。
lambdaBias	DOUBLE	ラムダ バイアス線型ブースティングパラメータ。任意の数値を指定します。デフォルトは 0 です。
numBoostRound バージョン 18.2.1.1 以降は num_boost_round に名前変更されています。	integer	モデル作成用の num boost round 値。1 から 1000 の間の値を指定します。デフォルトは 10 です。

表 270. *xgboostlinearnode* のプロパティ (続き)

xgboostlinearnode のプロパティ	データ・タイプ	プロパティの説明
objectiveType	string	学習タスクの目的タイプ。指定できる値は <code>reg:linear</code> 、 <code>reg:logistic</code> 、 <code>reg:gamma</code> 、 <code>reg:tweedie</code> 、 <code>count:poisson</code> 、 <code>rank:pairwise</code> 、 <code>binary:logistic</code> 、または <code>multi</code> です。フラグ型対象の場合、 <code>binary:logistic</code> または <code>multi</code> のみを使用できます。 <code>multi</code> を使用する場合、スコア結果には XGBoost 目的タイプ <code>multi:softmax</code> および <code>multi:softprob</code> が表示されます。
random_seed	integer	乱数シード。0 から 9999999 の範囲の任意の数値です。デフォルトは 0 です。
useHPO	Boolean	<code>true</code> または <code>false</code> に指定して、HPO オプションを有効または無効にします。 <code>true</code> に設定すると <code>Rbfopt</code> が適用され、自動的に「最適な」One-Class SVM モデルが判別されます。この場合、ユーザーが <code>target_objval</code> パラメータで定義した目標値に到達します。

xgboosttreenode のプロパティ



XGBoost Tree® は、ツリーモデルを基本モデルとして使用する勾配ブースティング・アルゴリズムの高度な実装です。ブースティング・アルゴリズムでは、弱い分類子に繰り返し学習させ、それを最終的な強い分類子に追加します。XGBoost ツリーは柔軟性が極めて高く、多くのユーザーを圧倒するほどの多数のパラメータが用意されています。このため、スポンサーモデラーの XGBoost ツリー・ノードでは、コア・フィーチャーおよびよく使用されるパラメータが公開されています。このノードは Python で実装されています。

表 271. *xgboosttreenode* のプロパティ

xgboosttreenode のプロパティ	データ・タイプ	プロパティの説明
TargetField バージョン 18.2.1.1 以降は <code>target</code> に名前変更されています。	フィールド	対象フィールド。
InputFields バージョン 18.2.1.1 以降は <code>inputs</code> に名前変更されています。	フィールド	入力フィールド。
treeMethod バージョン 18.2.1.1 以降は <code>tree_method</code> に名前変更されています。	string	モデルの構築用のツリー手法。指定可能な値は、 <code>auto</code> 、 <code>exact</code> 、または <code>approx</code> です。デフォルトは <code>auto</code> です。

表 271. *xgboosttreenode* のプロパティ (続き)

xgboosttreenode のプロパティ	データ・タイプ	プロパティの説明
numBoostRound バージョン 18.2.1.1 以降は num_boost_round に名前変更されています。	<i>integer</i>	モデル作成用の num boost round 値。1 から 1000 の間の値を指定します。デフォルトは 10 です。
maxDepth バージョン 18.2.1.1 以降は max_depth に名前変更されています。	<i>integer</i>	ツリーの成長の最大深度。1 以上の値を指定します。デフォルトは 6 です。
minChildWeight バージョン 18.2.1.1 以降は min_child_weight に名前変更されています。	<i>DOUBLE</i>	ツリーの成長のための子の重みの最小値。0 以上の値を指定します。デフォルトは 1 です。
maxDeltaStep バージョン 18.2.1.1 以降は max_delta_step に名前変更されています。	<i>DOUBLE</i>	ツリーの成長の差分ステップの最大数。0 以上の値を指定します。デフォルトは 0 です。
objectiveType バージョン 18.2.1.1 以降は objective_type に名前変更されています。	<i>string</i>	学習タスクの目的タイプ。指定できる値は reg:linear、reg:logistic、reg:gamma、reg:tweedie、count:poisson、rank:pairwise、binary:logistic、または multi です。フラグ型対象の場合、binary:logistic または multi のみを使用できます。multi を使用する場合、スコア結果には XGBoost 目的タイプ multi:softmax および multi:softprob が表示されます。
earlyStopping バージョン 18.2.1.1 以降は early_stopping に名前変更されています。	<i>Boolean</i>	早期停止機能を使用するかどうか。デフォルトは False です。
earlyStoppingRounds バージョン 18.2.1.1 以降は early_stopping_rounds に名前変更されています。	<i>integer</i>	学習を続行するには、少なくとも早期停止ラウンドごとに検証エラーが減少する必要があります。デフォルトは 10 です。
evaluationDataRatio バージョン 18.2.1.1 以降は evaluation_data_ratio に名前変更されています。	<i>DOUBLE</i>	検証エラーに使用する入力データの比率。デフォルトは 0.3 です。

表 271. `xgboosttreenode` のプロパティ (続き)

xgboosttreenode のプロパティ	データ・タイプ	プロパティの説明
<code>random_seed</code>	<i>integer</i>	乱数シード。0 から 9999999 の範囲の任意の数値です。デフォルトは 0 です。
<code>sampleSize</code> バージョン 18.2.1.1 以降は <code>sample_size</code> に名前変更されています。	<i>DOUBLE</i>	過剰適合を制御するためのサブサンプル。0.1 から 1.0 の間の値を指定します。デフォルトは 0.1 です。
<code>eta</code>	<i>DOUBLE</i>	過剰適合を制御するためのイータ。0 から 1 の間の値を指定します。デフォルトは 0.3 です。
<code>gamma</code>	<i>DOUBLE</i>	過剰適合を制御するためのガンマ。0 以上の任意の数値を指定してください。デフォルトは 6 です。
<code>colsSampleRatio</code> バージョン 18.2.1.1 以降は <code>col_sample_ratio</code> に名前変更されています。	<i>DOUBLE</i>	過剰適合を制御するためのツリー別の列サンプル。0.01 から 1 の間の値を指定します。デフォルトは 1 です。
<code>colsSampleLevel</code> バージョン 18.2.1.1 以降は <code>col_sample_level</code> に名前変更されています。	<i>DOUBLE</i>	過剰適合を制御するためのレベル別の列サンプル。0.01 から 1 の間の値を指定します。デフォルトは 1 です。
<code>lambda</code>	<i>DOUBLE</i>	過剰適合を制御するためのラムダ。0 以上の任意の数値を指定してください。デフォルトは 1 です。
<code>alpha</code>	<i>DOUBLE</i>	過剰適合を制御するためのアルファ。0 以上の任意の数値を指定してください。デフォルトは 0 です。
<code>scalePosWeight</code> バージョン 18.2.1.1 以降は <code>scale_pos_weight</code> に名前変更されています。	<i>DOUBLE</i>	不均衡データ・セットを扱うためのスケールの正の重み。デフォルトは 1 です。
<code>use_HP0</code> バージョン 18.2.1.1 で追加されました。		

第 20 章 Spark ノードのプロパティ

isotonicasnode プロパティ



Isotonic 回帰は、回帰アルゴリズムのファミリーに属します。スpos モデラーの Isotonic-AS ノードは Spark で実装されています。Isotonic 回帰アルゴリズムについて詳しくは、<https://spark.apache.org/docs/2.2.0/mllib-isotonic-regression.html> を参照してください。

表 272. *isotonicasnode* プロパティ

isotonicasnode プロパティ	データ・タイプ	プロパティの説明
label	<i>string</i>	このプロパティは、Isotonic 回帰が計算される対象の従属変数です。
features	<i>string</i>	このプロパティは、独立変数です。
weightCol	<i>string</i>	重みは、測度の数を表します。デフォルトは 1 です。
isotonic	<i>Boolean</i>	このプロパティは、タイプが <i>isotonic</i> または <i>antitonic</i> のどちらであることを示します。
featureIndex	<i>integer</i>	このプロパティは、 <i>featuresCol</i> がベクトル列である場合の機能のインデックス用です。デフォルトは 0 です。

kmeansasnode プロパティ



K-Means は、最も一般的に使用されるクラスタリングアルゴリズムの 1 つです。このアルゴリズムは、データポイントをクラスタリングして、事前定義された数のクラスタを作成します。スpos モデラーの K-Means-AS ノードは Spark で実装されています。K-Means アルゴリズムについて詳しくは、<https://spark.apache.org/docs/2.2.0/ml-clustering.html> を参照してください。K-Means-AS ノードでは、カテゴリ変数の場合にワンホットエンコーディングが自動的に実行されることに留意してください。

表 273. *kmeansasnode* プロパティ

kmeansasnode プロパティ	値	プロパティの説明
roleUse	<i>string</i>	定義済みの役割を使用するには <i>predefined</i> を指定し、ユーザー設定フィールドの割り当てを使用するには <i>custom</i> を指定します。デフォルトは <i>predefined</i> です。

表 273. kmeansasnode プロパティ (続き)

kmeansasnode プロパティ	値	プロパティの説明
autoModel	Boolean	新しく生成されたスコアリング フィールドにデフォルト名 (\$S-prediction) を使用するには true を指定し、カスタム名を使用するには false を指定します。 デフォルトは true です。
features	フィールド	roleUse プロパティが custom に設定されている場合の入力用のフィールド名のリスト。
name	string	autoModel プロパティが false に設定されている場合の、新しく生成されたスコアリング フィールドの名前。
clustersNum	integer	作成するクラスターの数。 デフォルトは 5 です。
initMode	string	初期化アルゴリズム。 指定できる値は、k-means または random です。 デフォルトは k-means です。
initSteps	integer	initMode が k-means に設定されている場合の初期化ステップの数。 デフォルトは 2 です。
advancedSettings	Boolean	次の 4 つのプロパティを使用可能にするには true を指定します。 デフォルトは false です。
maxIteration	integer	クラスタリングの最大反復数。 デフォルトは 20 です。
tolerance	string	反復を停止する許容度。 可能な設定は、1.0E-1、1.0E-2、...、1.0E-6 です。 デフォルトは 1.0E-4 です。
setSeed	Boolean	カスタム ランダム シードを使用するには true を指定します。 デフォルトは false です。
randomSeed	integer	setSeed プロパティが true の場合のカスタム ランダム シード。

multilayerperceptronnode のプロパティ



多層パーセプトロンはフィード フォワード人工ニューラル ネットワークに基づく分類器であり、複数の層から構成されます。 各層は、ネットワーク内の次の層に全結合されます。 スポス モデラー の MultiLayerPerceptron-AS ノードは Spark で実装されています。 多層パーセプトロン分類器 (MLPC) について詳しくは、<https://spark.apache.org/docs/latest/ml-classification-regression.html#multilayer-perceptron-classifier> を参照してください。

表 274. *multilayerperceptronnode* のプロパティ

multilayerperceptronnode のプロパティ	データ・タイプ	プロパティの説明
features	フィールド	予測の入力として使用する 1 つ以上のフィールド。
label	フィールド	予測の目標として使用するフィールド。
layers[0]	<i>integer</i>	含めるパーセプトロン層の数。デフォルトは 1 です。
layers[1...<latest-1>]	<i>integer</i>	隠れ層の数。デフォルトは 1 です。
layers[<latest>]	<i>integer</i>	出力層の数。デフォルトは 1 です。
seed	<i>integer</i>	カスタム ランダム シード。
maxiter	<i>integer</i>	実行する反復の最大回数。デフォルトは 10 です。

xgboostasnode プロパティ



XGBoost は、勾配ブースティング・アルゴリズムの高度な実装です。ブースティング・アルゴリズムでは、弱い分類子に繰り返し学習させ、それを最終的な強い分類子に追加します。XGBoost は柔軟性が極めて高く、多くのユーザーを圧倒するほどの多数のパラメータが用意されています。このため、スポンサー モデラー の XGBoost-AS ノードでは、コア・フィーチャーおよびよく使用されるパラメータが公開されています。XGBoost-AS ノードは Spark で実装されています。

表 275. *xgboostasnode* プロパティ

xgboostasnode プロパティ	データ・タイプ	プロパティの説明
target_field	フィールド	対象のフィールド名のリスト。
input_fields	フィールド	入力用のフィールド名のリスト。
nWorkers	<i>integer</i>	XGBoost モデルの学習に使用するワーカーの数。デフォルトは 1 です。
numThreadPerTask	<i>integer</i>	ワーカーごとに使用するスレッドの数。デフォルトは 1 です。
useExternalMemory	<i>Boolean</i>	キャッシュとして外部メモリーを使用するかどうか。デフォルトは False です。
boosterType	<i>string</i>	使用するブースティング タイプ。選択可能なオプションは <i>gbtree</i> 、 <i>gblinear</i> 、または <i>dart</i> です。デフォルトは <i>gbtree</i> です。
numBoostRound	<i>integer</i>	ブースティングのラウンド数。0 以上の値を指定します。デフォルトは 10 です。
scalePosWeight	<i>DOUBLE</i>	正の重みと負の重みのバランスを制御します。デフォルトは 1 です。
randomseed	<i>integer</i>	乱数発生ルーチンによって使用されるシード。デフォルトは 0 です。

表 275. xgboostasnode プロパティ (続き)

xgboostasnode プロパティ	データ・タイプ	プロパティの説明
objectiveType	string	学習目的。指定できる値は reg:linear、reg:logistic、reg:gamma、reg:tweedie、rank:pairwise、binary:logistic、または multi です。フラグ型対象の場合、binary:logistic または multi のみを使用できます。multi を使用する場合、スコア結果には XGBoost 目的タイプ multi:softmax および multi:softprob が表示されます。デフォルトは reg:linear です。
evalMetric	string	検証データの評価メトリック。デフォルトのメトリックは、目的に応じて割り当てられます。指定できる値は rmse、mae、logloss、error、merror、mlogloss、auc、ndcg、map、または gamma-deviance です。デフォルトは rmse です。
lambda	DOUBLE	重みに関する L2 正規化項。この値を大きくすると、モデルがより保守的になります。0 以上の任意の数値を指定してください。デフォルトは 1 です。
alpha	DOUBLE	重みに関する L1 正規化項。この値を大きくすると、モデルがより保守的になります。0 以上の任意の数値を指定してください。デフォルトは 0 です。
lambdaBias	DOUBLE	偏りに関する L2 正規化項。gblinear ブースティングタイプが使用されている場合は、このラムダ バイアス線型ブースティングパラメータを使用できます。0 以上の任意の数値を指定してください。デフォルトは 0 です。
treeMethod	string	gbtree または dart ブースティングタイプが使用されている場合は、ツリーの成長のためのこのツリー方法パラメータ (および以降のその他のツリーパラメータ) を使用できます。これは、使用する XGBoost ツリー構築アルゴリズムを指定します。使用可能なオプションは、auto、exact、または approx です。デフォルトは auto です。
maxDepth	integer	ツリーの最大深度。2 以上の値を指定します。デフォルトは 6 です。
minChildWeight	DOUBLE	子に必要なインスタンスの重み (ヘシアン) の最小合計。0 以上の値を指定します。デフォルトは 1 です。

表 275. xgboostasnode プロパティ (続き)

xgboostasnode プロパティ	データ・タイプ	プロパティの説明
maxDeltaStep	DOUBLE	各ツリーの重みを推定できるようにするための差分ステップの最大数。0 以上の値を指定します。デフォルトは 0 です。
sampleSize	DOUBLE	サブサンプルは、学習インスタンスの比率を示します。0.1 から 1.0 の間の値を指定します。デフォルトは 1.0 です。
eta	DOUBLE	オーバーフィッティングを防ぐために更新ステップ中に使用するステップサイズの収縮。0 から 1 の間の値を指定します。デフォルトは 0.3 です。
gamma	DOUBLE	ツリーの葉ノードをさらに分割するために必要な最小の損失低減。0 以上の任意の数値を指定してください。デフォルトは 6 です。
colsSampleRatio	DOUBLE	各ツリーを構築する際の、列のサブサンプルの比率。0.01 から 1 の間の値を指定します。デフォルトは 1 です。
colsSampleLevel	DOUBLE	各分割における、各レベルでの列のサブサンプルの比率。0.01 から 1 の間の値を指定します。デフォルトは 1 です。
normalizeType	string	dart ブースティング タイプが使用されている場合は、この dart パラメータおよび以下の 3 つの dart パラメータを使用できます。このパラメータは、正規化アルゴリズムを設定します。tree または forest を指定します。デフォルトは tree です。
sampleType	string	サンプリング アルゴリズム タイプ。uniform または weighted を指定します。デフォルトは uniform です。
rateDrop	DOUBLE	ドロップアウト率 dart ブースティング パラメータ。0.0 から 1.0 の間の値を指定します。デフォルトは 0.0 です。
skipDrop	DOUBLE	スキップ ドロップアウトの確率の dart ブースティング パラメータ。0.0 から 1.0 の間の値を指定します。デフォルトは 0.0 です。

第 21 章 スーパーノードのプロパティ

スーパーノード固有のプロパティを次の表に示します。共通のノード・プロパティもスーパーノードに適用されることに注意してください。

表 276. ターミナル・スーパーノードのプロパティ		
プロパティ名	プロパティの種類/値のリスト	プロパティの説明
execute_method	Script	
	Normal	
script	string	

スーパーノードのパラメーター

次の一般形式を使用して、スーパーノードのパラメーターを作成または設定するためにスクリプトを使用できます。

```
mySuperNode.setParameterValue("minvalue", 30)
```

以下を使用して、パラメーター値を取得することができます。

```
value mySuperNode.getParameterValue("minvalue")
```

既存のスーパーノードの検索

findByType() 関数を使用して、ストリーム内のスーパーノードを検索できます。

```
source_supernode = modeler.script.stream().findByType("source_super", None)
process_supernode = modeler.script.stream().findByType("process_super", None)
terminal_supernode = modeler.script.stream().findByType("terminal_super", None)
```

カプセル化ノードのプロパティ設定

スーパーノード内の子ダイアグラムにアクセスすることにより、スーパーノードの中にカプセル化された特定のノードのプロパティを設定できます。例えば、データを読み込むためにカプセル化された可変長ファイルのある入力スーパーノードがあるとします。以下のようにして、子ダイアグラムにアクセスし、関連ノードを検索することにより、読み込みファイルの名前 (full_filename プロパティを使用して指定) を渡すことができます。

```
childDiagram = source_supernode.getChildDiagram()
varfilenode = childDiagram.findByType("variablefile", None)
varfilenode.setPropertyValue("full_filename", "c:/mydata.txt")
```

スーパーノードの作成

スーパーノードとその中身を初めから作成する場合、同様の方法で行うことができます。このためには、スーパーノードを作成し、子ダイアグラムにアクセスして、目的のノードを作成します。スーパーノード

のダイアグラム内のすべてのノードを、入力コネクター・ノードや出力コネクター・ノードとリンクさせるようにすることも必要です。例えば、プロセス スーパーノードを作成する場合は、次のようにします。

```
process_supernode = modeler.script.stream().createAt("process_super", "My  
SuperNode", 200, 200)  
childDiagram = process_supernode.getChildDiagram()  
filternode = childDiagram.createAt("filter", "My Filter", 100, 100)  
childDiagram.linkFromInputConnector(filternode)  
childDiagram.linkToOutputConnector(filternode)
```

付録 A ノード名のリファレンス

ここでは、IBM SPSS Modeler のノードのスクリプト名のリファレンスを提供します。

モデル・ナゲット名

モデル・ナゲット (生成されたモデル) は、ノード・オブジェクトと出力オブジェクトと同様に、その種類で参照できます。次の表に、モデル・オブジェクトの参照名を一覧表示します。

これらの名前は、IBM SPSS Modeler ウィンドウの右上隅にある「モデル」パレット内のモデル・ナゲットを参照するために、特に使用されます。スコアリングの目的でストリームに追加されたモデル・ノードを参照するには、`apply...` の接頭辞が付いた別の名前セットが使用されます。

注: 通常の場合では、名前および種類の両方でモデルを参照することが、混乱を避けるために推奨されます。

表 277. モデル・ナゲット名 (「モデル作成」パレット)	
モデル名	モデル
<code>anomalydetection</code>	異常値
<code>Apriori</code>	Apriori
<code>autoclassifier</code>	自動分類
<code>autocluster</code>	自動クラスタリング
<code>autonumeric</code>	自動数値
<code>bayesnet</code>	Bayesian network (ベイズ)
<code>c50</code>	C5.0
<code>carma</code>	Carma
<code>cart</code>	C&R Tree
<code>chaid</code>	CHAID
<code>coxreg</code>	Cox 回帰
<code>decisionlist</code>	ディシジョン・リスト
<code>discriminant</code>	判別分析
<code>factor</code>	因子分析
<code>featureselection</code>	特徴量選択
<code>genlin</code>	一般化線型
<code>glmm</code>	GLMM
<code>kmeans</code>	K-Means
<code>knn</code>	<i>k</i> 最近傍法
<code>kohonen</code>	Kohonen
線型	線形
<code>logreg</code>	ロジスティック回帰
<code>neuralnetwork</code>	ニューラル・ネットワーク

表 277. モデル・ナゲット名 (「モデル作成」パレット) (続き)

モデル名	モデル
quest	QUEST
線型回帰	線型回帰
シーケンス	シーケンス
slrm	自己学習応答モデル
statisticsmodel	IBM SPSS Statistics モデル
svm	Support Vector Machine
timeseries	時系列
TwoStep	TwoStep

表 278. モデル・ナゲット名 (「データベース・モデリング」パレット)

モデル名	モデル
db2imcluster	IBM ISW クラスタリング
db2imlog	IBM ISW ロジスティック回帰
db2imnb	IBM ISW Naive Bayes
db2imreg	IBM ISW 回帰
db2imtree	IBM ISW ディシジョン・ツリー
msassoc	MS アソシエーション・ルール
msbayes	MS Naive Bayes
mscluster	MS クラスタリング
mslogistic	MS Logistic Regression
msneuralnetwork	MS ニューラル・ネットワーク
msregression	MS Linear Regression
mssequencecluster	MS Sequence Clustering
mstimeseries	MS Time Series
mstree	MS ディシジョン・ツリー
netezzabayes	Netezza ベイズ・ネットワーク
netezzadectree	Netezza ディシジョン・ツリー
netezzadivcluster	Netezza 分裂クラスタリング
netezzaglm	Netezza 一般化線型
netezzakmeans	Netezza K-Means
netezzaknn	Netezza KNN
netezzalinereregression	Netezza 線型
netezzanaivebayes	Netezza Naive Bayes
netezzapca	Netezza PCA
netezzaregtree	Netezza 回帰ツリー

表 278. モデル・ナゲット名 (「データベース・モデリング」パレット) (続き)	
モデル名	モデル
netezzatimeseries	Netezza 時系列
oraabn	Oracle Adaptive Bayes
oraai	Oracle AI
oradecisiontree	Oracle ディシジョン・ツリー
oraglm	Oracle GLM
orakmeans	Oracle k-Means
oranb	Oracle Naive Bayes
oranmf	Oracle NMF
oraocluster	Oracle O-Cluster
orasvm	Oracle SVM

重複するモデル名の回避

生成されたモデルを操作するのにスクリプトを使用する場合、重複するモデル名を使用していると、スクリプトがあいまいになることに注意する必要があります。これを避けるために、スクリプト作成時に、生成されたモデルには一意の名前を使用することをお勧めします。

重複するモデル名に関するオプションを設定するには

1. メニューから次の項目を選択します。
「ツール」 > 「ユーザー オプション」
2. 「通知」 タブをクリックします。
3. 生成されたモデルに対して重複する名前を禁止するには、「**前のモデルを置換**」を選択します。

あいまいなモデルの参照がある場合、スクリプト実行の動作は SPSS Modeler と IBM SPSS Collaboration and Deployment Services との間で異なります。SPSS Modeler クライアントには自動的に同じ名前を持つモデルを置き換えるオプション「以前のモデルを置き換える」があります (例えば、スクリプトをループで反復して随時異なる名前を作成)。しかし、このオプションは、同じスクリプトが IBM SPSS Collaboration and Deployment Services で実行される場合は使用できません。ループの終了前に、モデルに対するあいまいな参照を回避するために各反復で生成されるモデルの名前を変更するか、現在のモデルをクリアすることにより (clear generated palette 文の追加など)、この状況を回避することができます。

出力形式名

次の表に、すべての出力オブジェクトの形式と、それを作成するノードを一覧表示します。

表 279. 出力オブジェクトの種類と、そのオブジェクトを作成するノード	
出力オブジェクトの種類	ノード
analysisoutput	分析
collectionoutput	コレクション
dataauditoutput	データ検査
distributionoutput	分布
evaluationoutput	評価
histogramoutput	ヒストグラム

表 279. 出力オブジェクトの種類と、そのオブジェクトを作成するノード (続き)	
出力オブジェクトの種類	ノード
matrixoutput	クロス集計
meansoutput	平均値
multiplotoutput	マルチ散布図
plotoutput	散布図
qualityoutput	品質
reportdocumentoutput	このオブジェクトの種類はノードからのものではなく、プロジェクト・レポートに作成された出力です。
reportoutput	報告書
statisticsprocedureoutput	統計 出力
statisticsoutput	統計
tableoutput	表
timeplotoutput	時系列グラフ
weboutput	Web

付録 B 従来のスクリプトから Python スクリプトへの移行

従来のスクリプトの移行の概要

ここでは、IBM SPSS Modeler での Python スクリプトと従来のスクリプトの違いを要約し、従来のスクリプトを Python スクリプトに移行する方法について説明します。また、SPSS Modeler の標準的な従来のコマンドと、同等の Python コマンドのリストも示します。

一般的な差異

従来のスクリプトの設計の大部分は、OS コマンド・スクリプトが基になっています。従来のスクリプトは、行指向であり、一部のブロック構造 (if...then...else...endif や for...endfor など) があるとしても、インデントには一般に意味がありません。

Python スクリプトでは、インデントには意味があり、同一の論理ブロックに属する複数の行は、同じレベルにインデントされている必要があります。

注: Python コードをコピーして貼り付ける場合は、注意が必要です。タブを使用してインデントされている行は、エディター上で、スペースを使用してインデントされている行と同じように見える場合があります。しかし、これらの行が同じインデントであるとは見なされないため、Python スクリプトはエラーを生成します。

スクリプト・コンテキスト

スクリプト・コンテキストは、スクリプトを実行する環境 (例えば、スクリプトを実行するストリームやスーパーノード) を定義します。従来のスクリプトでは、コンテキストは暗黙的です。つまり、例えば、ストリーム・スクリプト内のノード参照は、そのスクリプトを実行するストリーム内にあると想定されます。

Python スクリプトでは、スクリプト・コンテキストは、`modeler.script` モジュールによって明示的に提供されます。例えば、Python ストリーム・スクリプトは、以下のコードを使用して、スクリプトを実行するストリームにアクセスできます。

```
s = modeler.script.stream()
```

ストリームに関連した関数は、返されたオブジェクトによって呼び出すことができます。

コマンドと関数

従来のスクリプトは、コマンド指向です。つまり、スクリプトの各行は、実行する必要があるコマンドが先頭にあり、パラメーターが後に続きます。例えば、以下のとおりです。

```
connect 'Type':typenode to :filternode  
rename :derivnode as "Compute Total"
```

Python は、通常、関数を定義するオブジェクト (モジュール、クラス、またはオブジェクト) によって起動される関数を使用します。例えば、以下のとおりです。

```
stream = modeler.script.stream()  
typenode = stream.findByType("type", "Type")  
filternode = stream.findByType("filter", None)  
stream.link(typenode, filternode)  
derive.setLabel("Compute Total")
```

リテラルとコメント

IBM SPSS Modeler でよく使用される一部のリテラル・コマンドおよびコメント・コマンドには、Python スクリプトの同等コマンドがあります。これは、SPSS Modeler の既存の従来のスクリプトを、IBM SPSS Modeler 17 で使用できるように、Python スクリプトに変換するのに役立ちます。

表 280. リテラルとコメントの従来のスクリプトから Python スクリプトへのマッピング	
従来のスクリプト	Python スクリプト
整数。例: 4	Same
浮動小数点数。例: 0.003	Same
単一引用符で囲まれた文字列。例: 'Hello'	Same 注: 非 ASCII 文字が含まれている文字列リテラルには、接頭辞 u を付けて、Unicode として表します。
二重引用符で囲まれた文字列。例: "Hello again"	Same 注: 非 ASCII 文字が含まれている文字列リテラルには、接頭辞 u を付けて、Unicode として表します。
長い文字列。例: """This is a string that spans multiple lines"""	Same
リスト。例: [1 2 3]	[1, 2, 3]
変数の参照。例: set x = 3	x = 3
行の継続 (\)。例: set x = [1 2 \ 3 4]	x = [1, 2,\ 3, 4]
ブロックのコメント。例: /* This is a long comment over a line. */	""" This is a long comment over a line. """
行のコメント。例: set x = 3 # make x 3	x = 3 # make x 3
undef	None
true	True
false	False

演算子

IBM SPSS Modeler でよく使用される一部の演算子コマンドには、Python スクリプトの同等コマンドがあります。これは、SPSS Modeler の既存の従来のスクリプトを、IBM SPSS Modeler 17 で使用できるように、Python スクリプトに変換するのに役立ちます。

表 281. 演算子の従来のスクリプトから Python スクリプトへのマッピング

従来のスクリプト	Python スクリプト
NUM1 + NUM2 LIST + ITEM LIST1 + LIST2	NUM1 + NUM2 LIST.append(ITEM) LIST1.extend(LIST2)
NUM1 - NUM2 LIST - ITEM	NUM1 - NUM2 LIST.remove(ITEM)
NUM1 * NUM2	NUM1 * NUM2
NUM1 / NUM2	NUM1 / NUM2
= ==	==
/= /==	!=
X ** Y	X ** Y
X < Y X <= Y X > Y X >= Y	X < Y X <= Y X > Y X >= Y
X div Y X rem Y X mod Y	X // Y X % Y X % Y
and or not(EXPR)	and or not EXPR

条件とループ

IBM SPSS Modeler でよく使用される一部の条件コマンドおよびループ・コマンドには、Python スクリプトの同等コマンドがあります。これは、SPSS Modeler の既存の従来のスクリプトを、IBM SPSS Modeler 17 で使用できるように、Python スクリプトに変換するのに役立ちます。

表 282. 条件とループの従来のスクリプトから Python スクリプトへのマッピング

従来のスクリプト	Python スクリプト
for VAR from INT1 to INT2 ... endfor	for VAR in range(INT1, INT2): ... または VAR = INT1 while VAR <= INT2: ... VAR += 1
for VAR in LIST ... endfor	for VAR in LIST: ...

表 282. 条件とループの従来のスクリプトから Python スクリプトへのマッピング (続き)

従来のスクリプト	Python スクリプト
<pre>for VAR in_fields_to NODE ... endfor</pre>	<pre>for VAR in NODE.getInputDataModel(): ...</pre>
<pre>for VAR in_fields_at NODE ... endfor</pre>	<pre>for VAR in NODE.getOutputDataModel(): ...</pre>
<pre>if...then ... elseif...then ... else ... endif</pre>	<pre>if ...: ... elif ...: ... else: ... </pre>
<pre>with TYPE OBJECT ... endwith</pre>	同等機能なし
var VAR1	変数宣言は不要

変数

従来のスクリプトでは、変数は参照される前に宣言します。例えば、以下のとおりです。

```
var mynode
set mynode = create typenode at 96 96
```

Python スクリプトでは、変数は初回の参照時に作成されます。例えば、以下のとおりです。

```
mynode = stream.createAt("type", "Type", 96, 96)
```

従来のスクリプトでは、変数の参照は ^ 演算子を使用して明示的に削除する必要があります。例えば、以下のとおりです。

```
var mynode
set mynode = create typenode at 96 96
set ^mynode.direction."Age" = Input
```

ほとんどのスクリプト言語と同様、Python スクリプトでは、これは不要です。例えば、以下のとおりです。

```
mynode = stream.createAt("type", "Type", 96, 96)
mynode.setKeyedPropertyValue("direction", "Age", "Input")
```

ノード、出力、およびモデルの各タイプ

従来のスクリプトのさまざまなオブジェクト・タイプ (ノード、出力、およびモデル) では、通常、タイプがオブジェクトのタイプに追加された形になっています。例えば、フィールド作成 (Derive) ノードのタイプは、`derivenode` です。

```
set feature_name_node = create derivenode at 96 96
```

Python の IBM SPSS Modeler API には、node 接尾辞が含まれないため、フィールド作成ノード (Derive) のタイプは、derive です。例えば、以下のとおりです。

```
feature_name_node = stream.createAt("derive", "Feature", 96, 96)
```

従来のスクリプトと Python スクリプトでのタイプ名の唯一の違いは、タイプ接尾辞がないことです。

プロパティ名

プロパティ名は、従来のスクリプトと Python スクリプトで同じです。例えば、可変長ファイル・ノードでは、ファイルの場所を定義するプロパティは、両方のスクリプト環境で `full_filename` です。

ノードの参照

多くの従来のスクリプトは、暗黙の検索を使用して、変更するノードを見つけてアクセスします。例えば、以下のコマンドは、ラベル「Type」を使用して、現行ストリームの中でデータ型ノードを検索し、「Age」フィールドの方向 (またはモデル作成の役割) を Input に、「Drug」フィールドを Target (予測される値) に設定します。

```
set 'Type':typenode.direction."Age" = Input
set 'Type':typenode.direction."Drug" = Target
```

Python スクリプトでは、プロパティ値を設定するための関数を呼び出す前に、ノード・オブジェクトを明示的に位置指定する必要があります。例えば、以下のとおりです。

```
typenode = stream.findByType("type", "Type")
typenode.setKeyedPropertyValue("direction", "Age", "Input")
typenode.setKeyedPropertyValue("direction", "Drug", "Target")
```

注: この場合、「Target」を文字列引用符で囲む必要があります。

Python スクリプトは、ModelingRole 列挙を `modeler.api` パッケージで 사용할 수도 있습니다。

Python スクリプトのバージョンは、より冗長な場合がありますが、ノードの検索は通常 1 回のみ行われるため、ランタイム・パフォーマンスが良くなります。従来のスクリプトの例では、ノードの検索は、コマンドごとに行われます。

ID によるノードの検索もサポートされています (ノード ID は、ノード・ダイアログの「注釈」タブで確認できます)。例えば、従来のスクリプトでは、以下のようになります。

```
# id65EMPB9VL87 is the ID of a Type node
set @id65EMPB9VL87.direction."Age" = Input
```

以下のスクリプトは、Python スクリプトを使用した場合の同じ例です。

```
typenode = stream.findByID("id65EMPB9VL87")
typenode.setKeyedPropertyValue("direction", "Age", "Input")
```

プロパティの取得と設定

従来のスクリプトは、`set` コマンドを使用して、値を割り当てます。`set` コマンドの後ろに、プロパティ定義を続けることができます。以下のスクリプトは、プロパティを設定するための 2 つの有効な形式を示しています。

```
set <node reference>.<property> = <value>
set <node reference>.<keyed-property>.<key> = <value>
```

Python スクリプトでは、関数 `setProperty()` と `setKeyedPropertyValue()` を使用して、同じ結果が得られます。例えば、以下のとおりです。

```
object.setPropertyValue(property, value)
object.setKeyedPropertyValue(keyed-property, key, value)
```

従来のスクリプトでは、`get` コマンドを使用して、プロパティ値にアクセスできます。例えば、以下のとおりです。

```
var n v
set n = get node :filternode
set v = ^n.name
```

Python スクリプトでは、関数 `getPropertyValue()` を使用して、同じ結果が得られます。例えば、以下のとおりです。

```
n = stream.findByType("filter", None)
v = n.getPropertyValue("name")
```

ストリームの編集

従来のスクリプトでは、`create` コマンドを使用して、新しいノードを作成します。例えば、以下のとおりです。

```
var agg select
set agg = create aggregatenode at 96 96
set select = create selectnode at 164 96
```

Python スクリプトでは、ノードを作成するためのさまざまなメソッドがストリームに用意されています。例えば、以下のとおりです。

```
stream = modeler.script.stream()
agg = stream.createAt("aggregate", "Aggregate", 96, 96)
select = stream.createAt("select", "Select", 164, 96)
```

従来のスクリプトでは、`connect` コマンドを使用して、ノード間のリンクを作成します。例えば、以下のとおりです。

```
connect ^agg to ^select
```

Python スクリプトでは、`link` メソッドを使用して、ノード間のリンクを作成します。例えば、以下のとおりです。

```
stream.link(agg, select)
```

従来のスクリプトでは、`disconnect` コマンドを使用して、ノード間のリンクを削除します。例えば、以下のとおりです。

```
disconnect ^agg from ^select
```

Python スクリプトでは、`unlink` メソッドを使用して、ノード間のリンクを削除します。例えば、以下のとおりです。

```
stream.unlink(agg, select)
```

従来のスクリプトでは、`position` コマンドを使用して、ストリーム・キャンバスにノードを配置したり、他のノード間にノードを配置したりします。例えば、以下のとおりです。

```
position ^agg at 256 256
position ^agg between ^myselect and ^mydistinct
```


Python スクリプトでは、2つの異なるメソッド `setXYPosition` と `setPositionBetween` を使用して、同じ結果が得られます。例：

```
agg.setXYPosition(256, 256)
agg.setPositionBetween(myselect, mydistinct)
```

ノード操作

IBM SPSS Modeler でよく使用される一部のノード操作コマンドには、Python スクリプトの同等コマンドがあります。これは、SPSS Modeler の既存の従来のスクリプトを、IBM SPSS Modeler 17 で使用できるように、Python スクリプトに変換するのに役立ちます。

表 283. ノード操作の従来のスクリプトから Python スクリプトへのマッピング

従来のスクリプト	Python スクリプト
<code>create nodespec at x y</code>	<pre>stream.create(type, name) stream.createAt(type, name, x, y) stream.createBetween(type, name, preNode, postNode) stream.createModelApplier(model, name)</pre>
<code>connect fromNode to toNode</code>	<code>stream.link(fromNode, toNode)</code>
<code>delete node</code>	<code>stream.delete(node)</code>
<code>disable node</code>	<code>stream.setEnabled(node, False)</code>
<code>enable node</code>	<code>stream.setEnabled(node, True)</code>
<code>disconnect fromNode from toNode</code>	<pre>stream.unlink(fromNode, toNode) stream.disconnect(node)</pre>
<code>duplicate node</code>	<code>node.duplicate()</code>
<code>execute node</code>	<pre>stream.runSelected(nodes, results) stream.runAll(results)</pre>
<code>flush node</code>	<code>node.flushCache()</code>
<code>position node at x y</code>	<code>node.setXYPosition(x, y)</code>
<code>position node between node1 and node2</code>	<code>node.setPositionBetween(node1, node2)</code>
<code>rename node as name</code>	<code>node.setLabel(name)</code>

ループ

従来のスクリプトでは、サポートされている主なループ・オプションが2つあります。

- カウント型 ループ。インデックス変数が、2つの整数の境界の間で変化します。
- シーケンス型 ループ。一連の値をループして、現在の値をループ変数にバインドします。

以下のスクリプトは、従来のスクリプトでのカウント型ループの例です。

```
for i from 1 to 10
  println ^i
endfor
```

以下のスクリプトは、従来のスクリプトでのシーケンス型ループの例です。

```
var items
set items = [a b c d]
```

```
for i in items
    println ^i
endfor
```

以下のような他のタイプのループも使用可能です。

- モデル・パレットのモデル、または出力パレットの出力を反復する。
- ノードに入るフィールドまたはノードから出るフィールドを反復する。

Python スクリプトでも、さまざまなタイプのループをサポートしています。以下のスクリプトは、Python スクリプトでのカウント型ループの例です。

```
i = 1
while i <= 10:
    print i
    i += 1
```

以下のスクリプトは、Python スクリプトでのシーケンス型ループの例です。

```
items = ["a", "b", "c", "d"]
for i in items:
    print i
```

シーケンス型ループは非常に柔軟であり、IBM SPSS Modeler API メソッドと組み合わせることにより、従来のスクリプトの大部分のユース・ケースをサポートできます。以下の例は、Python スクリプトでシーケンス型ループを使用して、ノードから出るフィールドを反復する方法を示しています。

```
node = modeler.script.stream().findByType("filter", None)
for column in node.getOutputDataModel().columnIterator():
    print column.getColumnname()
```

ストリームの実行

ストリームの実行中に、生成されたモデルまたは出力オブジェクトが、いずれかのオブジェクト・マネージャーに追加されます。従来のスクリプトでは、スクリプトは、作成されたオブジェクトをオブジェクト・マネージャーから位置指定するか、生成された最新の出力に、その出力を生成したノードからアクセスする必要があります。

Python でのストリームの実行は、実行により生成されたモデルまたは出力オブジェクトが、実行関数に渡されるリストに返されるという点で異なります。このため、ストリームの実行結果に、より簡単にアクセスできます。

従来のスクリプトは、以下の3つのストリーム実行コマンドをサポートしています。

- `execute_all` は、ストリーム内のすべての実行可能ターミナル・ノードを実行します。
- `execute_script` は、スクリプト実行の設定に関係なく、ストリーム・スクリプトを実行します。
- `execute node` は、指定したノードを実行します。

Python スクリプトは、以下のような同様の関数をサポートしています。

- `stream.runAll(results-list)` は、ストリーム内のすべての実行可能ターミナル・ノードを実行します。
- `stream.runScript(results-list)` は、スクリプト実行の設定に関係なく、ストリーム・スクリプトを実行します。
- `stream.runSelected(node-array, results-list)` は、指定したノードのセットを、指定した順に実行します。
- `node.run(results-list)` は、指定したノードを実行します。

従来のスクリプトでは、オプションの整数コードを指定した `exit` コマンドを使用して、ストリームの実行を終了できます。例えば、以下のとおりです。

```
exit 1
```

Python スクリプトでは、以下のスクリプトを使用して、同じ結果が得られます。

```
modeler.script.exit(1)
```

ファイル・システムおよびリポジトリによるオブジェクトへのアクセス

従来のスクリプトでは、`open` コマンドを使用して、既存のストリーム、モデル、または出力オブジェクトを開くことができます。例えば、以下のとおりです。

```
var s
set s = open stream "c:/my streams/modeling.str"
```

Python スクリプトには、セッションからアクセス可能で、同じような作業を実行できる `TaskRunner` クラスがあります。例えば、以下のとおりです。

```
taskrunner = modeler.script.session().getTaskRunner()
s = taskrunner.openStreamFromFile("c:/my streams/modeling.str", True)
```

従来のスクリプトを使用してオブジェクトを保存するには、`save` コマンドを使用します。例えば、以下のとおりです。

```
save stream s as "c:/my streams/new_modeling.str"
```

同等の Python スクリプトのアプローチでは、`TaskRunner` クラスを使用します。例えば、以下のとおりです。

```
taskrunner.saveStreamToFile(s, "c:/my streams/new_modeling.str")
```

IBM SPSS Collaboration and Deployment Services Repository ベースの操作は、`retrieve` および `store` コマンドを使用することによって、従来のスクリプトでサポートされています。例えば、以下のとおりです。

```
var s
set s = retrieve stream "/my repository folder/my_stream.str"
store stream ^s as "/my repository folder/my_stream_copy.str"
```

Python スクリプトでは、セッションに関連付けられているリポジトリ・オブジェクトによって、同等の機能にアクセスできます。

```
session = modeler.script.session()
repo = session.getRepository()
s = repo.retrieveStream("/my repository folder/my_stream.str", None, None, True)
repo.storeStream(s, "/my repository folder/my_stream_copy.str", None)
```

注：リポジトリにアクセスするには、有効なリポジトリ接続を使用してセッションが構成されている必要があります。

ストリーム操作

IBM SPSS Modeler でよく使用される一部のストリーム操作コマンドには、Python スクリプトの同等コマンドがあります。これは、SPSS Modeler の既存の従来のスクリプトを、IBM SPSS Modeler 17 で使用できるように、Python スクリプトに変換するのに役立ちます。

表 284. ストリーム操作の従来のスクリプトから Python スクリプトへのマッピング	
従来のスクリプト	Python スクリプト
create stream <i>DEFAULT_FILENAME</i>	<code>taskrunner.createStream(name, autoConnect, autoManage)</code>
close stream	<code>stream.close()</code>

表 284. ストリーム操作の従来のスクリプトから Python スクリプトへのマッピング (続き)	
従来のスクリプト	Python スクリプト
<code>clear stream</code>	<code>stream.clear()</code>
<code>get stream stream</code>	同等機能なし
<code>load stream path</code>	同等機能なし
<code>open stream path</code>	<code>taskrunner.openStreamFromFile(path, autoManage)</code>
<code>save stream as path</code>	<code>taskrunner.saveStreamToFile(stream, path)</code>
<code>retrieive stream path</code>	<code>repository.retrieiveStream(path, version, label, autoManage)</code>
<code>store stream as path</code>	<code>repository.storeStream(stream, path, label)</code>

モデルの操作

IBM SPSS Modeler でよく使用される一部のモデル操作コマンドには、Python スクリプトの同等コマンドがあります。これは、SPSS Modeler の既存の従来のスクリプトを、IBM SPSS Modeler 17 で使用できるように、Python スクリプトに変換するのに役立ちます。

表 285. モデル操作の従来のスクリプトから Python スクリプトへのマッピング	
従来のスクリプト	Python スクリプト
<code>open model path</code>	<code>taskrunner.openModelFromFile(path, autoManage)</code>
<code>save model as path</code>	<code>taskrunner.saveModelToFile(model, path)</code>
<code>retrieve model path</code>	<code>repository.retrieveModel(path, version, label, autoManage)</code>
<code>store model as path</code>	<code>repository.storeModel(model, path, label)</code>

ドキュメント出力操作

IBM SPSS Modeler でよく使用される一部のドキュメント出力操作コマンドには、Python スクリプトの同等コマンドがあります。これは、SPSS Modeler の既存の従来のスクリプトを、IBM SPSS Modeler 17 で使用できるように、Python スクリプトに変換するのに役立ちます。

表 286. ドキュメント出力操作の従来のスクリプトから Python スクリプトへのマッピング	
従来のスクリプト	Python スクリプト
<code>open output path</code>	<code>taskrunner.openDocumentFromFile(path, autoManage)</code>
<code>save output as path</code>	<code>taskrunner.saveDocumentToFile(output, path)</code>
<code>retrieve output path</code>	<code>repository.retrieveDocument(path, version, label, autoManage)</code>
<code>store output as path</code>	<code>repository.storeDocument(output, path, label)</code>

従来のスクリプトと Python スクリプトのその他の違い

レガシー・スクリプトは、IBM SPSS Modeler プロジェクトの操作をサポートしています。Python スクリプトは、現在、これをサポートしていません。

従来のスクリプトは、ステート型オブジェクト (ストリームおよびモデルの組み合わせ) をいくらかサポートしています。ステート型オブジェクトは、IBM SPSS Modeler 8.0 以降、廃止されました。Python スクリプトは、ステート型オブジェクトをサポートしていません。

Python スクリプトは、従来のスクリプトでは使用できない、以下の追加の機能を提供しています。

- クラス定義と関数定義
- エラー処理
- より高度な入出力サポート
- 外部のサード・パーティー・モジュール

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・ノード
プロパティ [221](#)
アソシエーション・ルール・ノード・ナゲット
のプロパティ [332](#)
暗号化パスワード
スクリプトへの追加 [54](#)
アンサンプル・ノード
プロパティ [163](#)
移行
アクセス、オブジェクトへの [465](#)
一般的な差異 [457](#)
概要 [457](#)
関数 [457](#)
コマンド [457](#)
出力タイプ [460](#)
スクリプト・コンテキスト [457](#)
ストリーム マネージャ、出力マネージャ、およびモデル
マネージャの消去 [35](#)
ストリームの実行 [464](#)
ストリームの編集 [462](#)
その他 [467](#)
ノード・タイプ [460](#)
ノードの参照 [461](#)
ファイル・システム [465](#)
プロパティの取得 [461](#)
プロパティの設定 [461](#)
プロパティ名 [461](#)
変数 [460](#)
モデルの種類 [460](#)
ループ [463](#)
repository [465](#)
異常値検出モデル
ノードのスクリプト・プロパティ [218](#), [331](#)
一般化線型モデル
ノードのスクリプト・プロパティ [253](#), [340](#)
因子分析モデル
ノードのスクリプト・プロパティ [249](#), [339](#)
エクスポート・ノード
ノードのスクリプト・プロパティ [403](#)
エラーのチェック
スクリプト [54](#)
オブジェクト指向 [24](#)

[カ行]

ガウス混合ノード
プロパティ [427](#), [432](#)
拡張のインポート・ノード
プロパティ [99](#)
拡張のエクスポート・ノード
プロパティ [413](#)
拡張の出力ノード

拡張の出力ノード (続き)
プロパティ [384](#)
拡張の変換ノード
プロパティ [129](#)
拡張モデル・ノード
ノードのスクリプト・プロパティ [246](#)
可変長ファイル・ノード
プロパティ [113](#)
関数
演算子 [458](#)
オブジェクト参照 [458](#)
コメント [458](#)
条件付き [459](#)
ストリーム操作 [465](#)
ドキュメント出力操作 [466](#)
ノード操作 [463](#)
モデルの操作 [466](#)
リテラル [458](#)
ループ [459](#)
記述統計ノード
プロパティ [395](#)
行列入替ノード
プロパティ [180](#)
クラスの作成 [25](#)
クラスの定義 [24](#)
グラフ作成ノード
スクリプトのプロパティ [191](#)
グラフボード・ノード
プロパティ [196](#)
グローバルの設定ノード
プロパティ [392](#)
クロス集計ノード
のプロパティ [386](#)
継承 [26](#)
構造化プロパティ [74](#)
コードのブロック [20](#)
固定長ファイル・ノード
プロパティ [101](#)
コマンド行
スクリプト [55](#)
引数のリスト [66](#), [68-70](#)
複数の引数 [70](#)
IBM SPSS Modeler の実行 [65](#)
parameters [67](#)

[サ行]

最近傍モデル
ノードのスクリプト・プロパティ [272](#)
再構成ノード
プロパティ [170](#)
再投影ノード
プロパティ [170](#)
座標系の再投影
プロパティ [170](#)
サポート・ベクター・マシン・モデル
ノードのスクリプト・プロパティ [347](#)

- サポート・ベクトル・マシン・モデル
 - ノードのスクリプト・プロパティ [307](#)
- 散布図ノード
 - プロパティ [208](#)
- サンプリング・ノード
 - プロパティ [135](#)
- シーケンス・モデル
 - ノードのスクリプト・プロパティ [297](#), [347](#)
- 時間区分ノード
 - プロパティ [174](#)
- 時間的因果モデル
 - ノードのスクリプト・プロパティ [308](#)
- 識別子 [19](#)
- 時空間予測ノード
 - プロパティ [300](#)
- 時系列グラフ・ノード
 - プロパティ [211](#)
- 時系列ノード
 - プロパティ [166](#)
- 時系列モデル
 - ノードのスクリプト・プロパティ [313](#), [321](#), [348](#), [349](#)
- 自己学習応答モデル
 - ノードのスクリプト・プロパティ [299](#), [347](#)
- システム
 - コマンド・ラインの引数 [66](#)
- 実行順序
 - スクリプトによる変更 [51](#)
- 自動クラスタリング・ノード
 - ノードのスクリプト・プロパティ [226](#)
- 自動クラスタリング・モデル
 - ノードのスクリプト・プロパティ [334](#)
- 自動数値モデル
 - ノードのスクリプト・プロパティ [228](#), [334](#)
- 自動分類ノード
 - ノードのスクリプト・プロパティ [224](#)
- 自動分類モデル
 - ノードのスクリプト・プロパティ [332](#)
- シミュレーション生成ノード
 - プロパティ [107](#)
- シミュレーションの当てはめノード
 - プロパティ [394](#)
- シミュレーション評価ノード
 - プロパティ [393](#)
- 集計棒グラフ・ノード
 - プロパティ [192](#)
- 主成分分析モデル
 - ノードのスクリプト・プロパティ [249](#), [339](#)
- 出力オブジェクト
 - スクリプト名 [455](#)
- 出力ノード
 - スクリプトのプロパティ [381](#)
- 順序ノード
 - プロパティ [169](#)
- 条件抽出ノード
 - プロパティ [137](#)
- 数学メソッド [21](#)
- 数値予測ノード・プロパティ [228](#)
- スーパーノード
 - スクリプト [1](#), [5](#), [6](#), [27](#), [451](#)
 - ストリーム [27](#)
 - パラメーター [451](#)
 - プロパティ [451](#)
 - プロパティの設定 [451](#)
 - streams [27](#)

- スクリプト
 - 以前のバージョンとの互換性 [55](#)
 - エラーのチェック [54](#)
 - 概要 [1](#), [15](#)
 - 共通のプロパティ [76](#)
 - グラフ作成ノード [191](#)
 - 構文 [15-17](#), [19-21](#), [23-26](#)
 - コマンド・ラインから [55](#)
 - コンテキスト [28](#)
 - 実行 [11](#)
 - 従来のスクリプト [458](#), [459](#), [463](#), [465](#), [466](#)
 - 出力ノード [381](#)
 - 条件付き実行 [6](#), [10](#)
 - 使用されている省略形 [74](#)
 - シンタックス [21](#)
 - スーパーノード・スクリプト [1](#), [27](#)
 - スーパーノード・ストリーム [27](#)
 - スーパーノード内 [5](#)
 - スタンドアロン スクリプト [1](#), [27](#)
 - ストリームの実行順序 [51](#)
 - ダイアグラム [27](#)
 - 中断 [11](#)
 - テキスト・ファイルからのインポート [1](#)
 - 特微量選択モデル [4](#)
 - 反復キー [8](#)
 - 反復変数 [9](#)
 - ビジュアル・ループ [6](#), [7](#)
 - フィールドの選択 [9](#)
 - 保存 [1](#)
 - ユーザーインターフェース [1](#), [4](#), [5](#)
 - ループ [6](#), [7](#)
 - Python スクリプト [458](#), [459](#), [463](#), [465](#), [466](#)
 - streams [1](#), [27](#)
- スクリプト API
 - エラーの処理 [44](#)
 - 概要 [39](#)
 - グローバル値 [48](#)
 - 検索 [39](#)
 - スーパーノードのパラメーター [44](#)
 - スタンドアロン スクリプト [49](#)
 - ストリーム・パラメーター [44](#)
 - 生成されたオブジェクトへのアクセス [42](#)
 - セッション・パラメーター [44](#)
 - ディレクトリーの取得 [39](#)
 - 複数ストリーム [49](#)
 - メタデータ [40](#)
 - 例 [39](#)
- スクリプトの実行 [11](#)
- スクリプトの中断 [11](#)
- スタンドアロン スクリプト [1](#), [4](#), [27](#)
- ステートメント [19](#)
- ストリーミング時系列分析ノード
 - プロパティ [146](#)
- ストリーミング時系列モデル
 - ノードのスクリプト・プロパティ [140](#)
- ストリーム実行結果へのアクセス
 - テーブル コンテンツ モデル [56](#)
 - JSON コンテンツ モデル [59](#)
 - XML コンテンツ モデル [58](#)
- ストリーム実行の結果へのアクセス
 - テーブル コンテンツ モデル [56](#)
 - JSON コンテンツ モデル [59](#)
 - XML コンテンツ モデル [58](#)
- ストリームでのループ [6](#), [7](#)

- ストリームの実行 [27](#)
- ストリームの実行順序
 - スクリプトによる変更 [51](#)
- ストリームの条件付き実行 [6](#), [10](#)
- ストリームの変更 [31](#), [34](#)
- ストリング
 - 大文字小文字の変換 [51](#)
- スペース・タイム・ボックス・ノード
 - プロパティ [126](#), [138](#)
- スペース・タイム・ボックス・ノードのプロパティ [126](#)
- スロット・パラメーター [5](#), [73](#), [75](#)
- 生成されたモデル
 - スクリプト名 [453](#), [455](#)
- 精度分析ノード
 - プロパティ [381](#)
- セキュリティ
 - 暗号化パスワード [54](#), [68](#)
- 線型回帰モデル
 - ノードのスクリプト・プロパティ [295](#), [346](#), [347](#)
- 線型サポート ベクター マシン モデル
 - ノードのスクリプト・プロパティ [284](#), [343](#)
- 線型モデル
 - ノードのスクリプト・プロパティ [276](#), [342](#)
- 線グラフ・ノード
 - プロパティ [207](#)
- 操作 [16](#)
- ソート・ノード
 - プロパティ [138](#)
- 属性の追加 [25](#)
- 属性の定義 [25](#)

[タ行]

- ダイアグラム [27](#)
- 置換ノード
 - プロパティ [164](#)
- 注釈 [19](#)
- 重複レコード・ノード
 - プロパティ [128](#)
- 地理空間入力ノード
 - プロパティ [106](#)
- ディシジョン・リスト・モデル
 - ノードのスクリプト・プロパティ [243](#), [337](#)
- データ型ノード
 - プロパティ [182](#)
- データ区分ノード
 - プロパティ [167](#)
- データ検査ノード
 - プロパティ [382](#)
- データの自動準備
 - プロパティ [152](#)
- データ分割ノード
 - プロパティ [157](#)
- データ分類ノード
 - プロパティ [168](#)
- データベース・エクスポート・ノード
 - プロパティ [406](#)
- データベース・ノード
 - プロパティ [93](#)
- データベース・モデル作成 [353](#)
- テーブル コンテンツ モデル [56](#)
- テーブル・ノード
 - プロパティ [396](#)
- 特微量選択モデル

- 特微量選択モデル (続き)
 - スクリプト [4](#)
 - 適用 [4](#)
 - ノードのスクリプト・プロパティ [251](#), [339](#)
- 匿名化ノード
 - プロパティ [151](#)

[ナ行]

- ナゲット
 - ノードのスクリプト・プロパティ [331](#)
- ニューラル・ネットワーク
 - ノードのスクリプト・プロパティ [288](#), [344](#)
- ニューラル・ネットワーク・モデル
 - ノードのスクリプト・プロパティ [285](#), [344](#)
- 入力ノード
 - のプロパティ [81](#)
- ノード
 - インポート [33](#)
 - 削除 [33](#)
 - 情報 [35](#)
 - スクリプトでのループ [51](#)
 - 置換 [33](#)
 - 名前のリファレンス [453](#)
 - ノードのリンク [32](#)
 - ノードのリンク解除 [32](#)
- ノードの検索 [29](#)
- ノードの作成 [31-33](#)
- ノードの参照
 - ノードの検索 [29](#)
 - プロパティの設定 [30](#)
- ノードのスクリプト・プロパティ
 - エクスポート・ノード [403](#)
 - モデル作成ノード [217](#)
 - モデル・ナゲット [331](#)
- ノードのトラバース [34](#)

[ハ行]

- パスワード
 - 暗号化 [68](#)
 - スクリプトへの追加 [54](#)
- パラメーター
 - スーパーノード [451](#)
 - スクリプト [15](#)
- バランス・ノード
 - プロパティ [122](#)
- 反復キー
 - スクリプトでのループ [8](#)
- 反復変数
 - スクリプトでのループ [9](#)
- 判別分析モデル
 - ノードのスクリプト・プロパティ [244](#), [337](#)
- 非 ASCII 文字 [23](#)
- 引数
 - コマンド・ファイル [70](#)
 - サーバー接続 [68](#)
 - システム [66](#)
 - IBM SPSS Analytic Server Repository 接続 [70](#)
 - IBM SPSS Collaboration and Deployment Services Repository の接続 [69](#)
- 引数の引き渡し [20](#)
- ヒストグラム・ノード

ヒストグラム・ノード (続き)

プロパティ [201](#)

非表示変数 [26](#)

評価ノード

プロパティ [194](#)

フィールド

スクリプトの無効化 [191](#)

フィールド作成ノード

プロパティ [160](#)

フィールド順序ノード

プロパティ [169](#)

フィールド名

大文字小文字の変換 [51](#)

フィルター・ノード

プロパティ [165](#)

フラグ設定ノード

プロパティ [172](#)

フラット・ファイル・ノード

プロパティ [414](#)

プロパティ

共通スクリプト [76](#)

スーパーノード [451](#)

スクリプト [73-75](#), [217](#), [331](#), [403](#)

ストリーム [77](#)

データベース・モデル作成ノード [353](#)

フィルター・ノード [74](#)

プロパティの設定 [30](#)

平均値ノード

プロパティ [389](#)

ベイズネット・プロパティ [230](#)

変換ノード

プロパティ [399](#)

変数

スクリプト [15](#)

棒グラフ・ノード

プロパティ [193](#)

[マ行]

マップ視覚化ノード

プロパティ [202](#)

メソッドの定義 [25](#)

文字列関数 [51](#)

モデル

スクリプト名 [453](#), [455](#)

モデル・オブジェクト

スクリプト名 [453](#), [455](#)

モデル作成ノード

ノードのスクリプト・プロパティ [217](#)

モデル・ナゲット

スクリプト名 [453](#), [455](#)

ノードのスクリプト・プロパティ [331](#)

[ヤ行]

ユーザー入力ノード

プロパティ [112](#)

[ラ行]

ランダム・フォレスト・ノード

プロパティ [435](#)

リスト [16](#)

ループ

スクリプトでの使用 [51](#)

例 [21](#)

レコード結合ノード

プロパティ [131](#)

レコード集計ノード

プロパティ [121](#)

レコード追加ノード

プロパティ [121](#)

レポート・ノード

プロパティ [391](#)

ロジスティック回帰モデル

ノードのスクリプト・プロパティ [279](#), [343](#)

A

aggregatenode プロパティ [121](#)

analysisnode プロパティ [381](#)

Analytic Server 入力ノード

プロパティ [88](#)

anomalydetectionnode プロパティ [218](#)

anonymizenode プロパティ [151](#)

appendnode プロパティ [121](#)

applyanomalydetectionnode プロパティ [331](#)

applyapriorinode プロパティ [331](#)

applyassociationrulesnode プロパティ [332](#)

applyautoclassifiernode プロパティ [332](#)

applyautoclusternode プロパティ [334](#)

applyautonumericnode プロパティ [334](#)

applybayesnetnode プロパティ [334](#)

applyc50node プロパティ [335](#)

applycarmanode プロパティ [335](#)

applycartnode プロパティ [335](#)

applychaidnode プロパティ [336](#)

applycoxregnode プロパティ [336](#)

applydecisionlistnode プロパティ [337](#)

applydiscriminantnode プロパティ [337](#)

applyextension プロパティ [338](#)

applyfactornode プロパティ [339](#)

applyfeatureselectionnode プロパティ [339](#)

applygeneralizedlinearnode プロパティ [340](#)

applygle プロパティ [341](#)

applyglmnode プロパティ [340](#)

applygmm のプロパティ [341](#)

applykmeansnode プロパティ [342](#)

applyknnnode プロパティ [342](#)

applykohonennode プロパティ [342](#)

applylinearnode プロパティ [343](#)

applylinearnode プロパティ [342](#)

applylogregnode プロパティ [343](#)

applylsvmnode プロパティ [343](#)

applymslogisticnode プロパティ [355](#)

applymsneuralnetworknode プロパティ [355](#)

applymsregressionnode プロパティ [355](#)

applymssequenceclusternode プロパティ [355](#)

applymstimeseriesnode プロパティ [355](#)

applymstreennode プロパティ [355](#)

applynetezabayesnode プロパティ [380](#)

applynetezadectreenode プロパティ [380](#)

applynetezadivclusternode プロパティ [380](#)

applynetezakmeansnode プロパティ [380](#)

applynetezaknnnode プロパティ [380](#)

applynetezalineregressionnode プロパティ [380](#)

applynetezanaivebayesnode プロパティ [380](#)

applynetezzapcanode プロパティ [380](#)
applynetezzaregtreenode プロパティ [380](#)
applyneuralnetnode プロパティ [344](#)
applyneuralnetworknode プロパティ [344](#)
applyocsvm のプロパティ [345](#)
applyoraabnnode プロパティ [364](#)
applyoradecisiontreenode プロパティ [364](#)
applyorakmeansnode プロパティ [364](#)
applyoranbnode プロパティ [364](#)
applyoranmfnode プロパティ [364](#)
applyoraoclusternode プロパティ [364](#)
applyorasvmnode プロパティ [364](#)
applyquestnode プロパティ [345](#)
applyr プロパティ [346](#)
applyrandomtrees プロパティ [346](#)
applyregressionnode プロパティ [347](#)
applyselflearningnode プロパティ [347](#)
applysequencenode プロパティ [347](#)
applystpnode プロパティ [348](#)
applysvmnode プロパティ [347](#)
appllytcnode プロパティ [348](#)
applytimeseriesnode プロパティ [349](#)
applytreeas プロパティ [349](#)
applyts プロパティ [348](#)
applytwestepAS のプロパティ [350](#)
applytwestepnode プロパティ [349](#)
applyxgboostlinearnode のプロパティ [350](#)
applyxgboosttreenode のプロパティ [350](#)
Apriori モデル
 ノードのスクリプト・プロパティ [219](#), [331](#)
apriorinode プロパティ [219](#)
AS 時間区分ノード
 プロパティ [156](#)
asexport プロパティ [403](#)
asimport プロパティ [88](#)
associationrulesnode プロパティ [221](#)
astimeintervalsnode プロパティ [156](#)
autoclassifiernode プロパティ [224](#)
autoclusternode プロパティ [226](#)
autodataprepnode プロパティ [152](#)
autonumericnode プロパティ [228](#)

B

balancenode プロパティ [122](#)
Bayesian network (ベイズ) モデル
 ノードのスクリプト・プロパティ [230](#), [334](#)
binningnode プロパティ [157](#)
buildr プロパティ [231](#)

C

C&R Tree モデル
 ノードのスクリプト・プロパティ [235](#), [335](#)
C5.0 モデル
 ノードのスクリプト・プロパティ [232](#), [335](#)
c50node プロパティ [232](#)
CARMA モデル
 ノードのスクリプト・プロパティ [234](#), [335](#)
carmanode プロパティ [234](#)
cartnode プロパティ [235](#)
CHAID モデル
 ノードのスクリプト・プロパティ [238](#), [336](#)

chaidnode プロパティ [238](#)
clear generated palette コマンド [55](#)
CLEM
 スクリプト [1](#)
cognosimport ノードのプロパティ [89](#)
collectionnode プロパティ [192](#)
Cox 回帰モデル
 ノードのスクリプト・プロパティ [240](#), [336](#)
coxregnode プロパティ [240](#)
CPLEX の最適化ノード
 プロパティ [123](#)
cplexoptnode プロパティ [123](#)

D

Data Collection エクスポート・ノード
 プロパティ [411](#)
Data Collection 入力ノード
 プロパティ [94](#)
dataauditnode プロパティ [382](#)
databaseexportnode プロパティ [406](#)
databasenode プロパティ [93](#)
datacollectionexportnode プロパティ [411](#)
datacollectionimportnode プロパティ [94](#)
decisionlist プロパティ [243](#)
derive_stbnode
 プロパティ [126](#)
derivinode プロパティ [160](#)
directedwebnode プロパティ [215](#)
discriminantnode プロパティ [244](#)
distinctnode プロパティ [128](#)
distributionnode プロパティ [193](#)

E

E-Plot ノード
 プロパティ [212](#)
ensemblenode プロパティ [163](#)
eplotnode プロパティ [212](#)
evaluationnode プロパティ [194](#)
Excel エクスポート・ノード
 プロパティ [412](#), [414](#)
Excel 入力ノード
 プロパティ [98](#)
excelexportnode プロパティ [412](#), [414](#)
excelimportnode プロパティ [98](#)
exportModelToFile [42](#)
extensionexportnode プロパティ [413](#)
extensionimportnode プロパティ [99](#)
extensionmodelnode プロパティ [246](#)
extensionoutputnode プロパティ [384](#)
extensionprocessnode プロパティ [129](#)

F

factornode プロパティ [249](#)
featureselectionnode プロパティ [4](#), [251](#)
fillernode プロパティ [164](#)
filternode プロパティ [165](#)
fixedfilenode プロパティ [101](#)
flags
 コマンド・ラインの引数 [65](#)
 複数のフラグの組み合わせ [70](#)

flatfilenode プロパティ [414](#)
for コマンド [51](#)

G

generated キーワード [55](#)
genlinnode プロパティ [253](#)
gle プロパティ [263](#)
GLE モデル
 ノードのスクリプト・プロパティ [263, 341](#)
GLMM モデル
 ノードのスクリプト・プロパティ [258, 340](#)
glmnode プロパティ [258](#)
gmm のプロパティ [427, 432](#)
graphboardnode プロパティ [196](#)
gsdata_import ノードのプロパティ [106](#)

H

HDBSCAN ノード
 プロパティ [428](#)
hdbscannode のプロパティ [428](#)
hdbscannugget のプロパティ [350](#)
histogramnode プロパティ [201](#)
historynode プロパティ [166](#)

I

IBM Cognos TM1 入力ノード
 プロパティ [109, 110](#)
IBM Cognos 入力ノード
 プロパティ [89](#)
IBM SPSS Analytic Server Repository
 コマンド・ラインの引数 [70](#)
IBM SPSS Collaboration and Deployment Services
 Repository
 コマンド・ラインの引数 [69](#)
 スクリプト [51](#)
IBM SPSS Modeler
 コマンド・ラインからの実行 [65](#)
IBM SPSS Statistics エクスポート・ノード
 プロパティ [425](#)
IBM SPSS Statistics 出力ノード
 プロパティ [425](#)
IBM SPSS Statistics 入力ノード
 プロパティ [423](#)
IBM SPSS Statistics 変換ノード
 プロパティ [423](#)
IBM SPSS Statistics モデル
 ノードのスクリプト・プロパティ [424](#)
Isotonic-AS ノード
 プロパティ [445](#)
isotonicasnode プロパティ [445](#)

J

JSON コンテンツ モデル [59](#)
JSON 入力ノード
 プロパティ [106](#)
jsonimportnode のプロパティ [106](#)
Jython [15](#)

K

K-Means モデル
 ノードのスクリプト・プロパティ [270, 342](#)
K-Means-AS モデル
 ノードのスクリプト・プロパティ [271, 445](#)
KDE シミュレーション・ノード
 プロパティ [385, 431](#)
KDE モデル
 ノードのスクリプト・プロパティ [351](#)
KDE モデル作成ノード
 プロパティ [429](#)
kdeapply のプロパティ [351](#)
kdeexport プロパティ [385, 431](#)
kdemodel のプロパティ [429](#)
kmeansasnode プロパティ [271, 445](#)
kmeansnode プロパティ [270](#)
KNN モデル
 ノードのスクリプト・プロパティ [342](#)
knnnode プロパティ [272](#)
Kohonen モデル
 ノードのスクリプト・プロパティ [274, 342](#)
kohonennode プロパティ [274](#)

L

linear プロパティ [276](#)
Linear-AS プロパティ [277](#)
Linear-AS モデル
 ノードのスクリプト・プロパティ [277, 343](#)
logregnode プロパティ [279](#)
lowertoupper 関数 [51](#)
LSVM モデル
 ノードのスクリプト・プロパティ [284](#)
lsvmnode プロパティ [284](#)

M

mapvisualization プロパティ [202](#)
matrixnode プロパティ [386](#)
meansnode プロパティ [389](#)
mergenode プロパティ [131](#)
Microsoft モデル
 ノードのスクリプト・プロパティ [353, 355](#)
MS Linear Regression
 ノードのスクリプト・プロパティ [353, 355](#)
MS Logistic Regression
 ノードのスクリプト・プロパティ [353, 355](#)
MS Sequence Clustering
 ノードのスクリプト・プロパティ [355](#)
MS Time Series
 ノードのスクリプト・プロパティ [355](#)
MS ディジション・ツリー
 ノードのスクリプト・プロパティ [353, 355](#)
MS ニューラル・ネットワーク
 ノードのスクリプト・プロパティ [353, 355](#)
msassocnode プロパティ [353](#)
msbayesnode プロパティ [353](#)
msclusternode プロパティ [353](#)
mslogisticnode プロパティ [353](#)
msneuralnetworknode プロパティ [353](#)
msregressionnode プロパティ [353](#)
mssequenceclusternode プロパティ [353](#)

mstimeseriesnode プロパティ [353](#)
mstreenode プロパティ [353](#)
MultiLayerPerceptron-AS ノード
プロパティ [446](#)
multilayerperceptronnode のプロパティ [446](#)
multiplotnode プロパティ [207](#)
multiset コマンド [74](#)

N

Netezza K-Means モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza KNN モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza Naive Bayes モデル
ノードのスクリプト・プロパティ [365](#)
Netezza Naive Bayesmodels
ノードのスクリプト・プロパティ [380](#)
Netezza 一般化線型モデル
ノードのスクリプト・プロパティ [365](#)
Netezza 回帰ツリー・モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza 時系列モデル
ノードのスクリプト・プロパティ [365](#)
Netezza 主成分分析モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza 線型モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza ディジション・ツリー・モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza 分裂クラスタリング・モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza ベイズ・ネットワーク・モデル
ノードのスクリプト・プロパティ [365, 380](#)
Netezza モデル
ノードのスクリプト・プロパティ [365](#)
netezabayesnode プロパティ [365](#)
netezadectreenode プロパティ [365](#)
netezadivclusternode プロパティ [365](#)
netezzaglmnode プロパティ [365](#)
netezzakmeansnode プロパティ [365](#)
netezzaknnnode プロパティ [365](#)
netezzalineressionnode プロパティ [365](#)
netezzanaivebayesnode プロパティ [365](#)
netezzapcanode プロパティ [365](#)
netezzaregtreenode プロパティ [365](#)
netezzatimeseriesnode プロパティ [365](#)
neuralnetnode プロパティ [285](#)
neuralnetworknode プロパティ [288](#)

O

ocsvmnode のプロパティ [433](#)
One-Class SVM ノード
プロパティ [433](#)
oraabnnode プロパティ [358](#)
oraainode プロパティ [358](#)
oraapriorinode プロパティ [358](#)
Oracle Adaptive Bayes モデル
ノードのスクリプト・プロパティ [358, 364](#)
Oracle AI モデル
ノードのスクリプト・プロパティ [358](#)
Oracle Apriori モデル

Oracle Apriori モデル (続き)
ノードのスクリプト・プロパティ [358, 364](#)
Oracle Decision Tree モデル
ノードのスクリプト・プロパティ [358, 364](#)
Oracle KMeans モデル
ノードのスクリプト・プロパティ [358, 364](#)
Oracle MDL モデル
ノードのスクリプト・プロパティ [358, 364](#)
Oracle Naive Bayes モデル
ノードのスクリプト・プロパティ [358, 364](#)
Oracle NMF モデル
ノードのスクリプト・プロパティ [358, 364](#)
Oracle O-Cluster
ノードのスクリプト・プロパティ [358, 364](#)
Oracle Support Vector Machines モデル
ノードのスクリプト・プロパティ [358, 364](#)
Oracle 一般化線型モデル
ノードのスクリプト・プロパティ [358](#)
Oracle モデル
ノードのスクリプト・プロパティ [358](#)
oradecisiontreenode プロパティ [358](#)
oraglmnode プロパティ [358](#)
orakmeansnode プロパティ [358](#)
oramdlnode プロパティ [358](#)
oranbnnode プロパティ [358](#)
oranmfnode プロパティ [358](#)
oraoclusternode プロパティ [358](#)
orasvmnode プロパティ [358](#)
outputfilenode プロパティ [414](#)

P

partitionnode プロパティ [167](#)
plotnode プロパティ [208](#)
Python
スクリプト [15](#)
Python モデル
ガウス混合ノードのスクリプト・プロパティ [341](#)
ノードのスクリプト・プロパティ [345, 350](#)

Q

QUEST モデル
ノードのスクリプト・プロパティ [290, 345](#)
questnode プロパティ [290](#)

R

R 構築ノード
ノードのスクリプト・プロパティ [231](#)
R 出力ノード
プロパティ [392](#)
R 変換ノード
プロパティ [134](#)
Random Trees モデル
ノードのスクリプト・プロパティ [293, 346](#)
randomtrees プロパティ [293](#)
reclassifynode プロパティ [168](#)
regressionnode プロパティ [295](#)
reordernode プロパティ [169](#)
reportnode プロパティ [391](#)
reprojectnode プロパティ [170](#)
restructurenode プロパティ [170](#)

retrieve コマンド [51](#)
RFM 集計ノード
 プロパティ [133](#)
RFM 分析ノード
 プロパティ [171](#)
rfmaggregatenode プロパティ [133](#)
rfanalysisnode プロパティ [171](#)
rfnode プロパティ [435](#)
routputnode のプロパティ [392](#)
Rprocessnode プロパティ [134](#)

S

samplenode プロパティ [135](#)
SAS エクスポート・ノード
 プロパティ [416](#)
SAS 入力ノード
 プロパティ [107](#)
sasexportnode プロパティ [416](#)
sasimportnode プロパティ [107](#)
selectnode プロパティ [137](#)
sequencenode プロパティ [297](#)
setglobalsnode プロパティ [392](#)
settoflagnode プロパティ [172](#)
simevalnode プロパティ [393](#)
simfitnode プロパティ [394](#)
simgenode プロパティ [107](#)
SLRM モデル
 ノードのスクリプト・プロパティ [299](#), [347](#)
slrmnode プロパティ [299](#)
SMOTE ノード
 プロパティ [437](#)
smotenode のプロパティ [437](#)
sortnode プロパティ [138](#)
spacetimeboxes プロパティ [138](#)
statisticsexportnode プロパティ [425](#)
statisticsimportnode プロパティ [4](#), [423](#)
statisticsmodelnode プロパティ [424](#)
statisticsnode プロパティ [395](#)
statisticsoutputnode プロパティ [425](#)
statisticstransformnode プロパティ [423](#)
store コマンド [51](#)
STP ノード
 プロパティ [300](#)
STP ノード ナゲット
 プロパティ [348](#)
stpnode プロパティ [300](#)
stream.nodes プロパティ [51](#)
streamingtimeseries プロパティ [140](#)
streamingts プロパティ [146](#)
streams
 実行 [27](#)
 条件付き実行 [6](#), [10](#)
 スクリプト [1](#), [27](#)
 プロパティ [77](#)
 変更 [31](#)
 ループ [6](#), [7](#)
 multiset コマンド [73](#)
SVM モデル
 ノードのスクリプト・プロパティ [307](#)
svmnode プロパティ [307](#)

T

t-SNE ノード
 プロパティ [213](#), [438](#)
tablenode プロパティ [396](#)
TCM モデル
 ノードのスクリプト・プロパティ [348](#)
tcmnode プロパティ [308](#)
timeintervalsnode プロパティ [174](#)
timeplotnode プロパティ [211](#)
timeseriesnode プロパティ [321](#)
tm1import ノードのプロパティ [110](#)
tm1odataimport ノードのプロパティ [109](#)
transformnode プロパティ [399](#)
transposenode プロパティ [180](#)
Tree-AS モデル
 ノードのスクリプト・プロパティ [324](#), [349](#)
treeas プロパティ [324](#)
ts プロパティ [313](#)
tsnode プロパティ [213](#), [438](#)
TWC インポート入力ノード
 プロパティ [111](#)
twcimport ノードのプロパティ [111](#)
TwoStep AS モデル
 ノードのスクリプト・プロパティ [327](#), [350](#)
TwoStep モデル
 ノードのスクリプト・プロパティ [326](#), [349](#)
twostepAS のプロパティ [327](#)
twostepnode プロパティ [326](#)
typenode プロパティ [4](#), [182](#)

U

userinputnode プロパティ [112](#)

V

variablefilenode プロパティ [113](#)

W

Web グラフ・ノード
 プロパティ [215](#)
webnode プロパティ [215](#)

X

XGBoost Linear ノード
 プロパティ [440](#)
XGBoost ツリー・ノード
 プロパティ [441](#)
XGBoost-AS ノード
 プロパティ [447](#)
xgboostasnode プロパティ [447](#)
xgboostlinearnode のプロパティ [440](#)
xgboosttreenode のプロパティ [441](#)
XML エクスポート・ノード
 プロパティ [421](#)
XML コンテンツ モデル [58](#)
XML 入力ノード
 プロパティ [118](#)
xmlexportnode プロパティ [421](#)
xmlimportnode プロパティ [118](#)

[特殊文字]

給仕人

コマンド・ラインの引数 [68](#)

