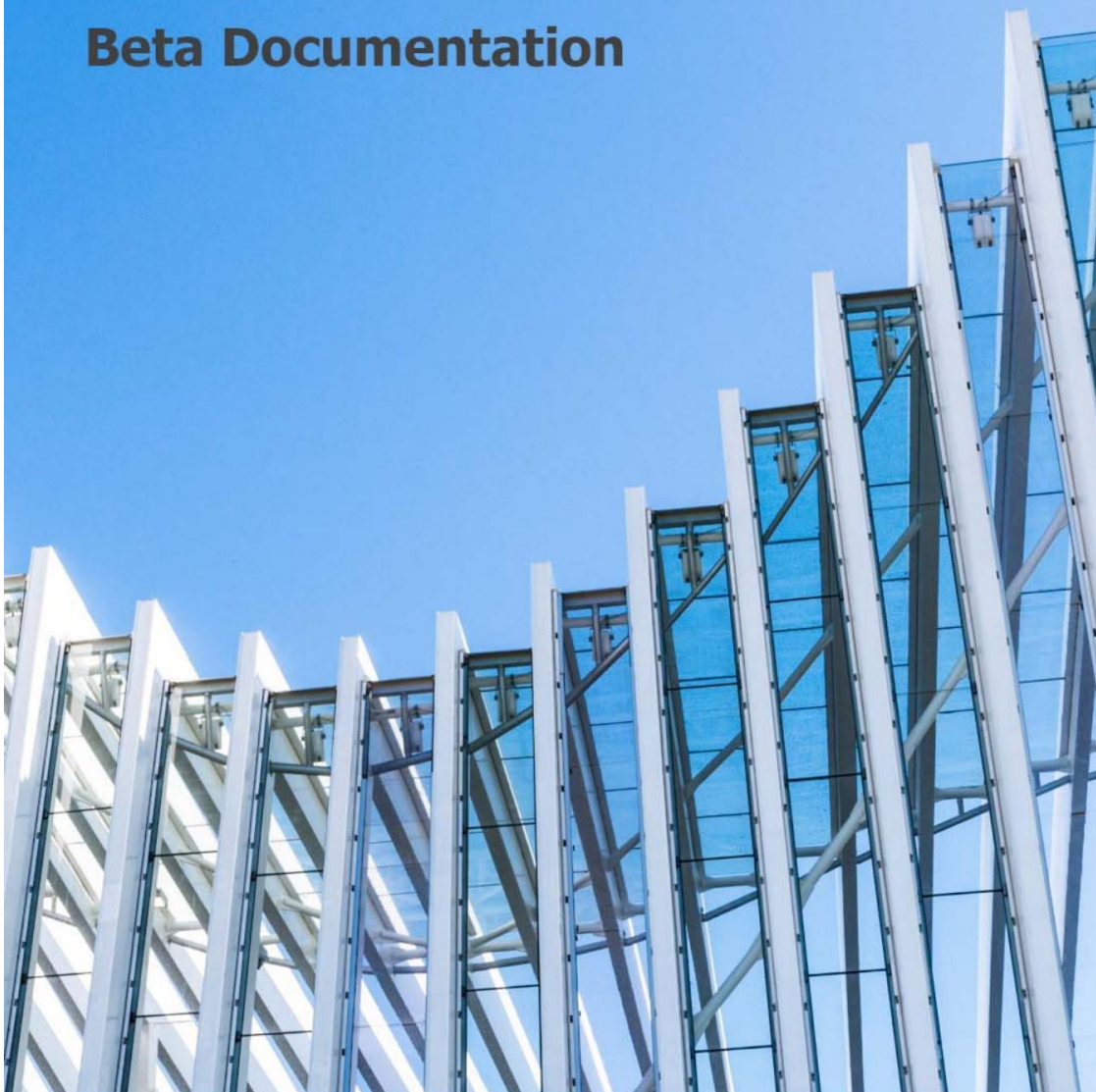


EViews[®] 15

Beta Documentation



EViews[®] 15 Beta 文档

版权所有 © 1994–2026 S&P Global Inc. 保留所有权利

2026 年 8 月 3 日

更多内容欢迎访问 www.eviews.com.cn

第 1 章 EViews 15 Beta 文档

EViews 开发团队很高兴地宣布，EViews 15 现已开放 Beta 测试。

除非您已在计算机上安装并注册了 EViews 14，否则 Beta 版无法运行。此外，本次发布的 Beta 版仅供有限时长使用，将在您运行程序时显示的到期日期之后停止运行。在 EViews 15 正式发布前，我们将持续为程序与文档提供定期更新和错误修复。

您可以通过以下电子邮件向我们发送评论、建议与错误报告：

support@eviews.com

如果您遇到任何问题或有任何建议，我们非常希望了解。报告问题时，请尽可能完整地描述问题的性质，包括您正在执行的操作序列以及出现的任何提示信息。请在往来邮件中附上您的 EViews 15 Beta 构建日期。您随时可以从 main menu 的 **Help/About EViews** 中查看您所用 EViews 副本的构建日期。

Beta 文档说明

本版文档的日期为 2026 年 8 月 3 日。

本文档包含 EViews 15 部分（并非全部）新功能的非常粗略的初步文档。

请注意，本文档部分内容摘自完整文档，因此部分内容可能格式不规范，索引不完整，且指向页面与章节的交叉引用链接可能失效。

尽管如此，您仍可以随意对文档提出意见；特别是，如果您发现某些部分不清楚或存在错误，请告知我们。

EViews 15 有哪些新特性？

EViews 15 带来了一系列令人振奋的变化与改进。以下列出了其中部分新功能。请注意，同一项功能可能出现在多个类别中。*这并非 EViews 15 新功能的完整列表。*

数据处理

季节调整

- 样条分解季节调整（“Spline Decomposition Seasonal Adjustment”，见第 3 页）。

计量经济学与统计

估计与分析

- 样条回归 (“Spline Regression”，见第 5 页)。
- 时间回归断点 (“Regression Discontinuity in Time”，见第 7 页)。
- 模型平均 (“Model Averaging”，见第 8 页)。
- 弹性网络 GLM (“Elastic-net GLM”，见第 14 页)。
- 因子增广 VAR (“Factor-Augmented VAR”，见第 15 页)。
- 贝叶斯 VAR 的增强 (“Enhancements to BVAR”，见第 15 页)。

预测

- 增强的 Facebook Prophet (“Enhanced Prophet”，见第 20 页)。
- NeuralProphet (“NeuralProphet”，见第 20 页)。
- 长短期记忆 (LSTM) 模型 (“LSTM Models”，见第 21 页)。
- AutoGluon (“AutoGluon”，见第 21 页)。
- Lag-Llama (“Lag-Llama”，见第 22 页)。
- 随机森林与决策树 (“Random Forests and Decision Trees”，见第 23 页)。
- Chronos-2 (“Chronos-2”，见第 24 页)。
- 增强的 ETS 平滑 (“Enhanced ETS Smoothing”，见第 24 页)。

命令参考的初步更新

- 新命令与已更新命令的初步列表 (第 4 章“命令参考的初步更新”，见第 25 页)。

第 2 章 数据处理

EViews 15 提供了一种新的季节调整方法。

季节调整

- 样条分解季节调整 (“Spline Decomposition Seasonal Adjustment”，见第 3 页)。

样条分解季节调整

EViews 15 新增了一种季节调整方法，能够使用样条回归对任意频率数据中的多个季节性成分进行建模。与传统季节调整方法相比，样条分解的优势在于它对时间序列的结构几乎不作假设，也无需严格的参数化。

EViews 允许为分解的趋势成分指定 B 样条 (B-spline)、P 样条 (P-spline) 或多项式样条 (polynomial spline) 之一，然后最多使用三个周期样条来建模季节成分。您还可以选择性地分解中额外添加线性回归变量，以建模节假日或一次性效应。

Spline Decomposition

Transformation
Method: Auto
Power: 2

Trend
Spline: p-Spline
Degree: 3
Penalty: 0.8

Seasonal

Season	Periodicity	Degree
☒ Season 1	12	3
☐ Season 2		
☐ Season 3		

Exogenous regressors (optional)

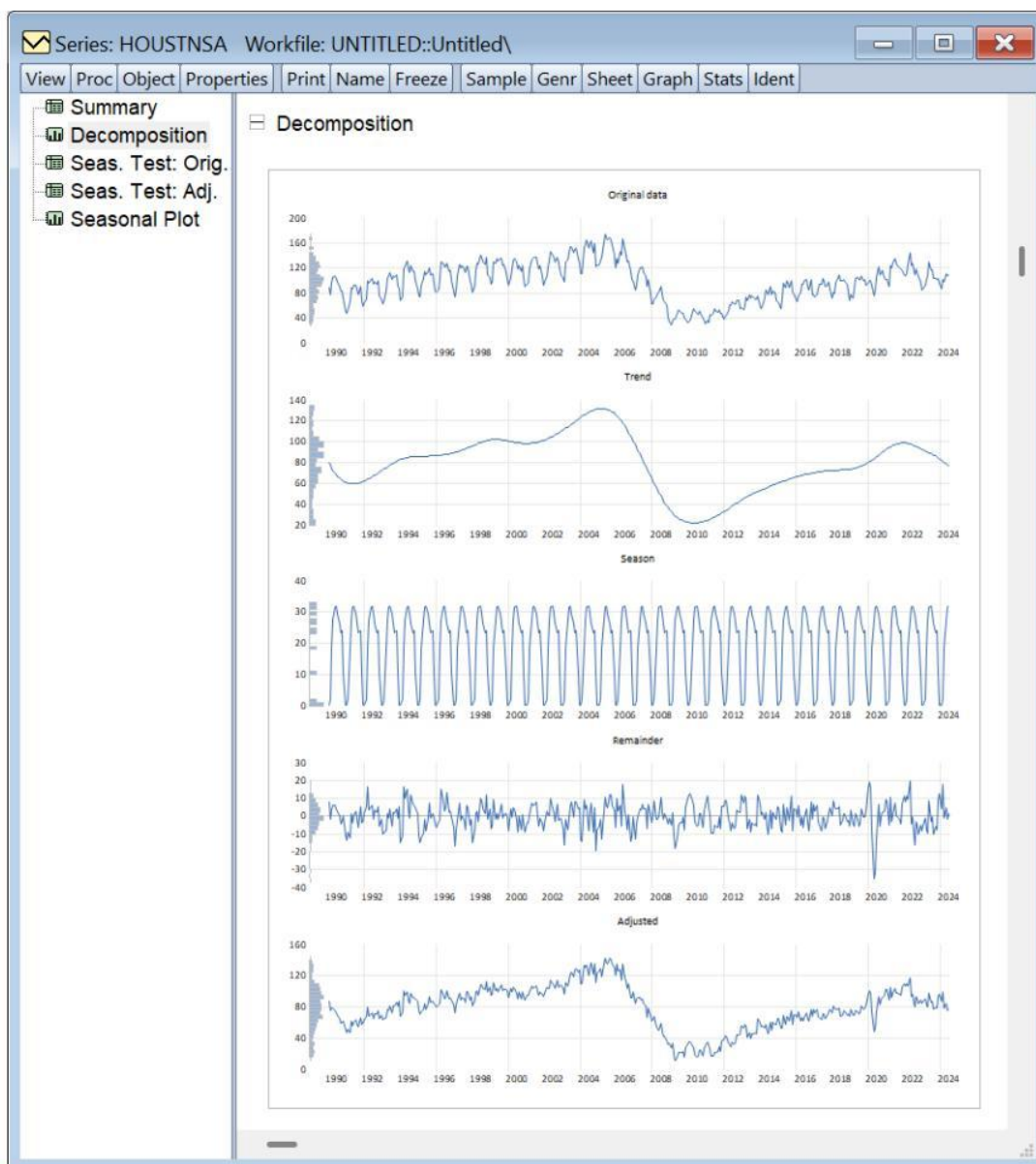
Sample specification
Decomposition sample: 1990m01 2024m06

Forecast sample (optional):
☒ Forecast length: 12
☐ Forecast endpoint:

Output
Season: houstnsa_saf
Trend: houstnsa_trend
Adjusted: houstnsa_sa

☒ Perform seasonality tests
☒ Display seasonality graphs

OK Cancel



新增或更新的命令

splinedecomp 使用样条分解对序列进行季节调整（第 74 页）。(新增)

第 3 章 计量经济学与统计

EViews 15 对其计量经济学与统计功能集进行了多项新增与改进。以下是最重要的新功能简介，随后是讨论与初步文档链接。

估计与分析

- 样条回归 (“Spline Regression”，见第 5 页)。
- 时间回归断点 (“Regression Discontinuity in Time”，见第 7 页)。
- 模型平均 (“Model Averaging”，见第 8 页)。
- 弹性网络 GLM (“Elastic-net GLM”，见第 14 页)。
- 因子增广 VAR (“Factor-Augmented VAR”，见第 15 页)。
- 贝叶斯 VAR 的增强 (“Enhancements to BVAR”，见第 15 页)。

预测

- 增强的 Facebook Prophet (“Enhanced Prophet”，见第 20 页)。
- NeuralProphet (“NeuralProphet”，见第 20 页)。
- 长短期记忆 (LSTM) 模型 (“LSTM Models”，见第 21 页)。
- AutoGluon (“AutoGluon”，见第 21 页)。
- Lag-Llama (“Lag-Llama”，见第 22 页)。
- 随机森林与决策树 (“Random Forests and Decision Trees”，见第 23 页)。
- Chronos-2 (“Chronos-2”，见第 24 页)。
- 增强的 ETS 平滑 (“Enhanced ETS Smoothing”，见第 24 页)。

样条回归

样条回归 (spline regression) 是一种通过拟合分段多项式曲线来建模非线性关系的强大技术，这些曲线在称为“节点 (knots)”的预定点上平滑连接。

EViews 15 引入了三种样条选项：

- B 样条：一组在有限子区间上非零的基函数，确保数值稳定性与局部控制。
- 多项式样条：在节点处连接、并对导数施加指定连续性约束的分段多项式。
- P 样条 (惩罚样条)：在相邻样条系数之差上附加粗糙度罚项以控制平滑程度的多项式样条。

EViews 还允许估计包含按线性处理的变量的样条回归。

Equation Estimation

Specification

Options

Equation specification

Dependent variable followed by list of non-parametric regressors

elec dmd tempf

List of linear regressors

Spline: bSpline

Degree: 3

Estimation settings

Method: SPLINEREG - Nonparametric spline estimation

Sample: 4/01/2005 4/30/2014

OK

Cancel

`partialeffects`.....查看使用样条估计的方程中偏效应的图形（第 66 页）。
(新增)

时间回归断点

时间回归断点（Regression Discontinuity in Time, RDIT）是一种准实验方法，用于估计在已知时间点发生的干预的因果效应。它将相对于干预日期的时间作为分配变量（running variable），并在该断点处寻找结果变量（outcome）的不连续跳跃，比较断点前后紧邻的观测值。其关键识别假设是：若没有干预，结果变量本应平滑地穿越断点演化，因此任何明显的突变都可归因于该处理（treatment）。

RDIT 的一个关键优势在于它只关注靠近干预断点的局部观测值，因此不像传统结构突变模型那样要求底层关系在整个样本期内保持稳定。

EViews 15 引入的 RDIT 遵循 Calonico、Cattaneo 与 Titiunik（2014a-c、2015）的局部多项式回归断点框架，及其向基于时间的分配变量的扩展。

新增或更新的命令

`rdit` 估计时间回归断点模型。（第 72 页）。(新增)

`balance` 查看使用协变量估计的时间回归断点模型的协变量平衡图（第 32 页）。
(新增)

`bwsensitivity` 查看时间回归断点模型的带宽敏感性图（第 35 页）。(新增)

`discontinuity` 查看时间回归断点模型的断点图（第 38 页）。(新增)

`donutsensitivity` 查看时间回归断点模型的甜甜圈敏感性图（第 38 页）。(新增)

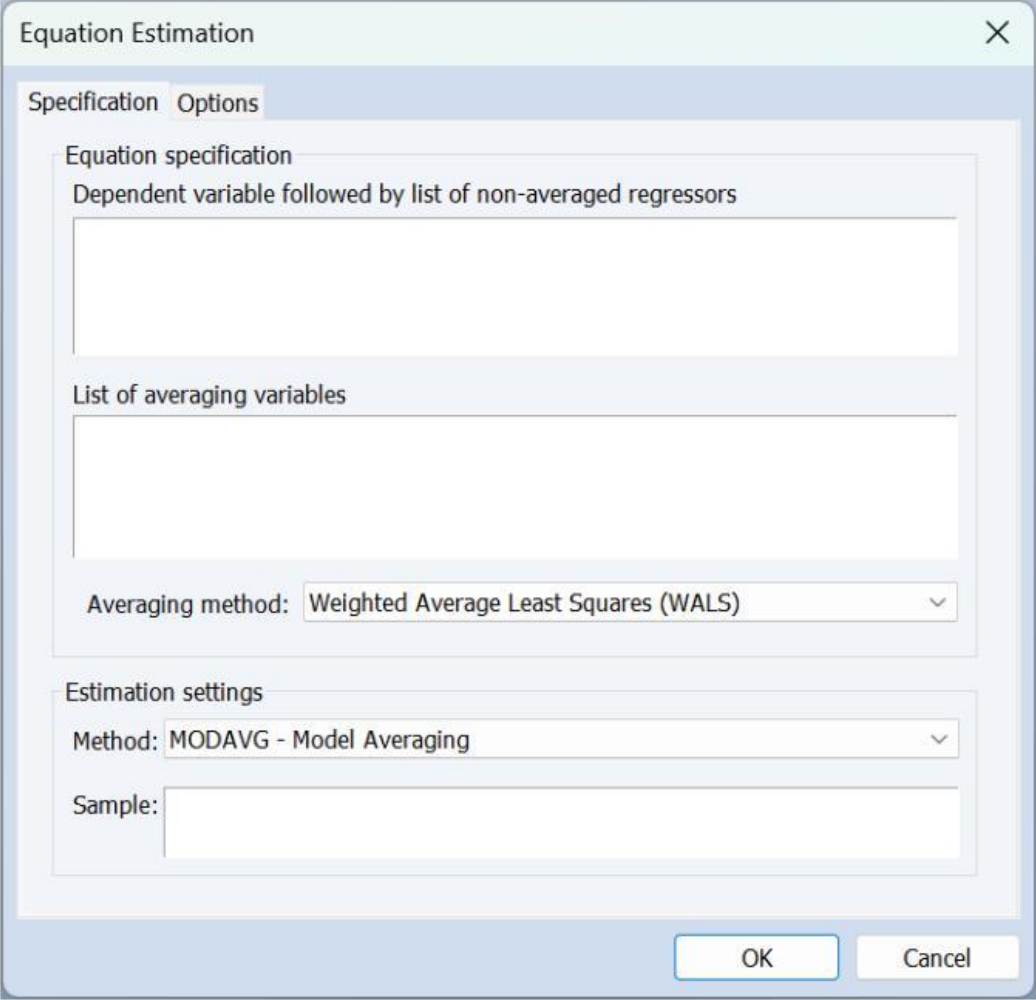
`placebo` 查看时间回归断点模型的安慰剂图（第 66 页）。(新增)

模型平均

模型平均是一种通过对多个模型的推断取平均来应对模型不确定性的估计技术。平均过程中使用模型权重，使得受数据支持的模型比不受支持的模型具有更大的影响。

EViews 15 提供了 Magnus 等人（2010）的模型平均方法：加权平均最小二乘（WALS）与贝叶斯模型平均（BMA）。WALS 是一种经典估计方法，使用近似后验模型概率的模型权重。BMA 是一种完全贝叶斯的技术，使用后验模型概率作为模型权重。有关该主题的深入讨论，请参见第 9 章“模型平均”，起始于第 289 页。

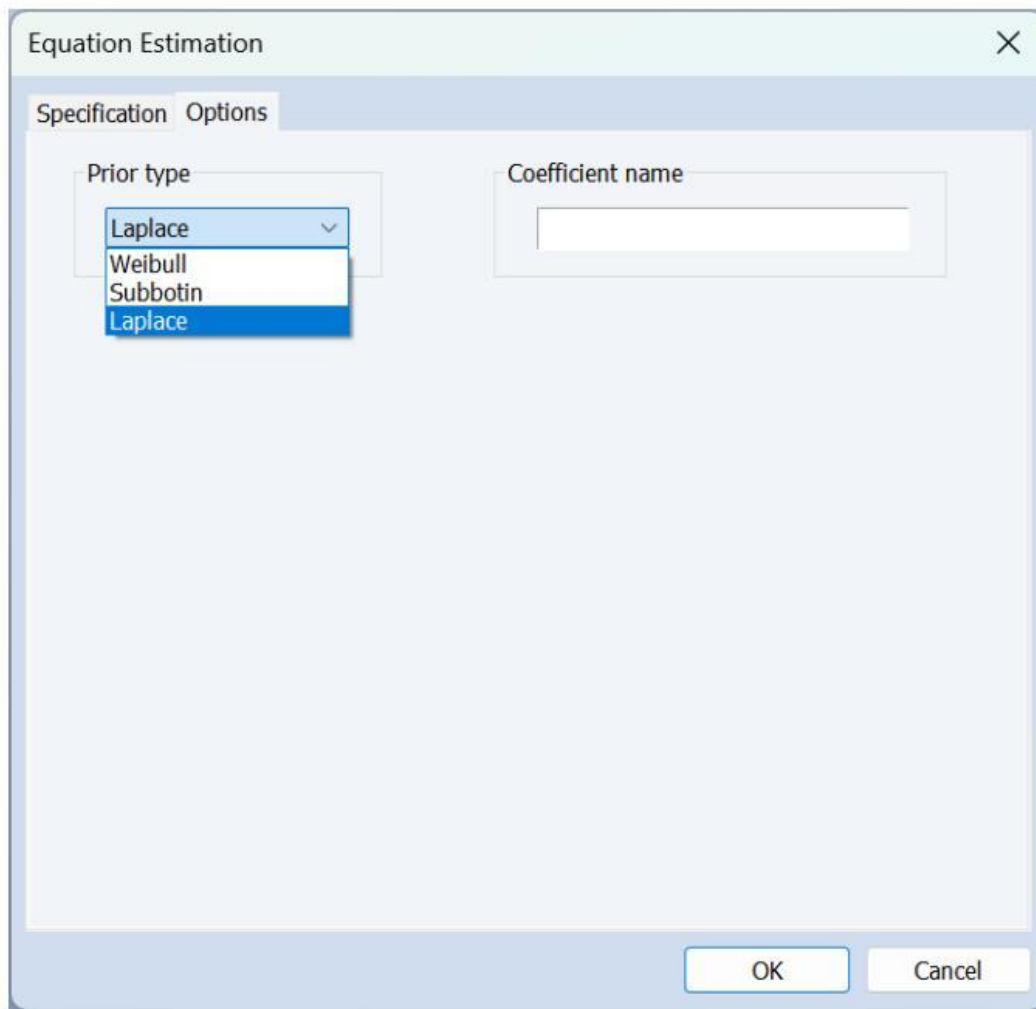
要在 EViews 中使用模型平均估计方程，请转到 **Quick/Estimate Equation...**，或在 command window 中输入 equation 并按 Enter。在 **Specification** 页面，从 **Method:** 下拉菜单中选择 **MODAVG - Model Averaging**。EViews 将显示如下对话框：



Equation specification 组框内有两个编辑字段。在第一个编辑字段中输入因变量名称，其后紧跟所需回归变量名称。在第二个编辑字段中输入可选回归变量的名称。

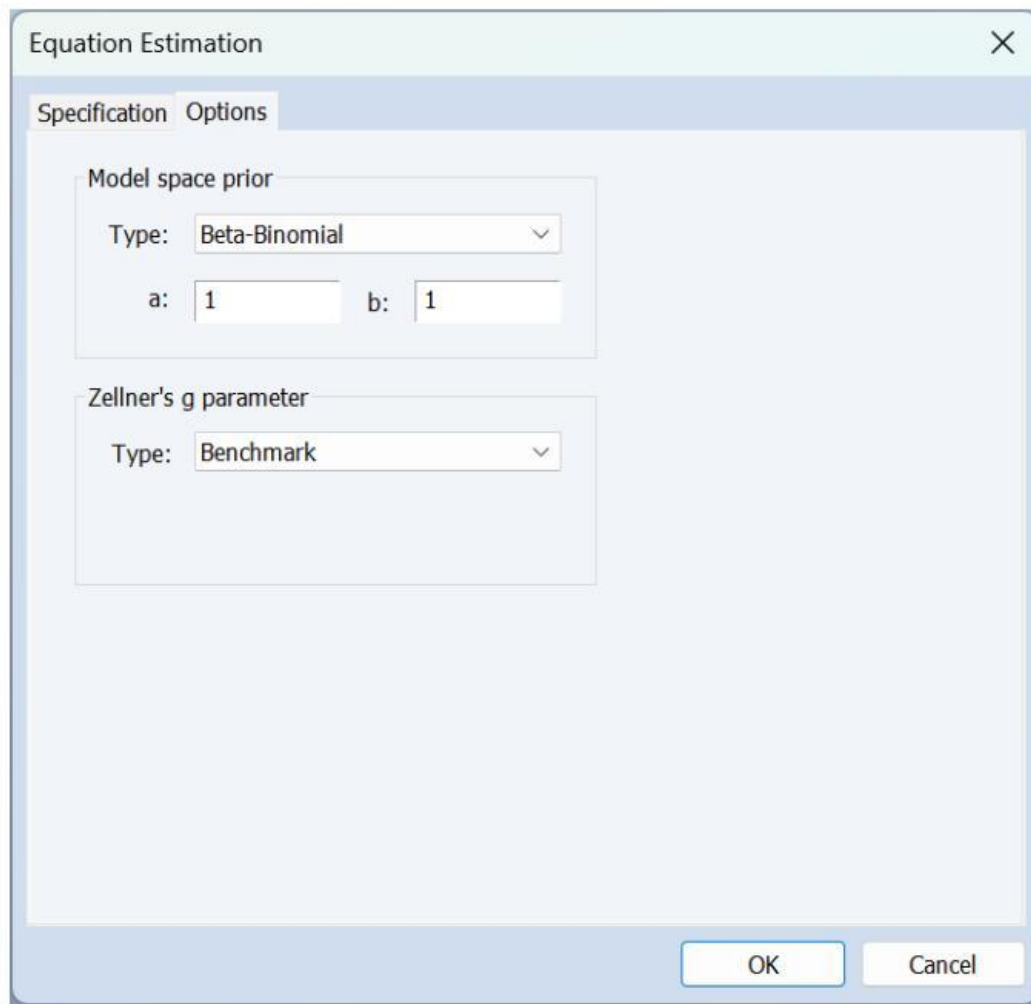
使用 **Averaging method:** 下拉菜单在 **Weighted Average Least Squares (WALS)** 与 **Bayesian Model Averaging (BMA)** 之间选择。所选择的平均方法将决定 **Options** 页面中可供您使用的选项。

若平均方法设置为 WALS，则 **Options** 页面允许您指定模型权重的近似分布：



详见“加权平均最小二乘（WALS）”，起始于第 298 页。

若平均方法设置为 BMA，则 **Options** 页面允许您指定模型空间与模型参数的先验分布：



详见“贝叶斯模型平均（BMA）”，起始于第 299 页。

在下文中， $\{M_\gamma\}$ 表示一组候选模型。

WALS

WALS 通过对来自 M_γ 的 OLS 估计量进行加权，形成平均估计量：

$$\hat{\beta}^{\text{WALS}} = \sum_{\gamma \in \{0, 1\}^{k_2}} w_\gamma \hat{\beta}_\gamma, \quad \sum_{\gamma} w_\gamma = 1, \quad w_\gamma \geq 0.$$

Magnus (2010) 指出，最优权重可通过将辅助回归变量（auxiliary regressors）相对于彼此以及相对于焦点回归变量（focus regressors）进行正交化来计算；详见“加权平均最小二乘（WALS）”，起始于第 298 页。

BMA

BMA 基于以下公式：

$$\pi(\Delta|y) = \sum_{\gamma} \pi(\Delta|y, M_{\gamma})\Pr(M_{\gamma}|y)$$

其中 $\pi(\Delta|y)$ 是某个感兴趣量 Δ 在所有候选模型上的边际后验分布， $\pi(\Delta|y, M_{\gamma})$ 是在给定模型 M_{γ} 条件下 Δ 的后验分布， $\Pr(M_{\gamma}|y)$ 为模型 M_{γ} 的后验模型概率。 Δ 的常见选择既包括共同参数 β 的集合，也包括预测值 y^* 。

Magnus 等人（2010）的 BMA 方法将上述公式应用于线性回归。更多细节，请参见“贝叶斯模型平均（BMA）”，起始于第 299 页。

示例

此处我们复现 Magnus 等人（2010）的 WALS 与 BMA 结果。完整示例请参见“示例”，起始于第 294 页。

WALS

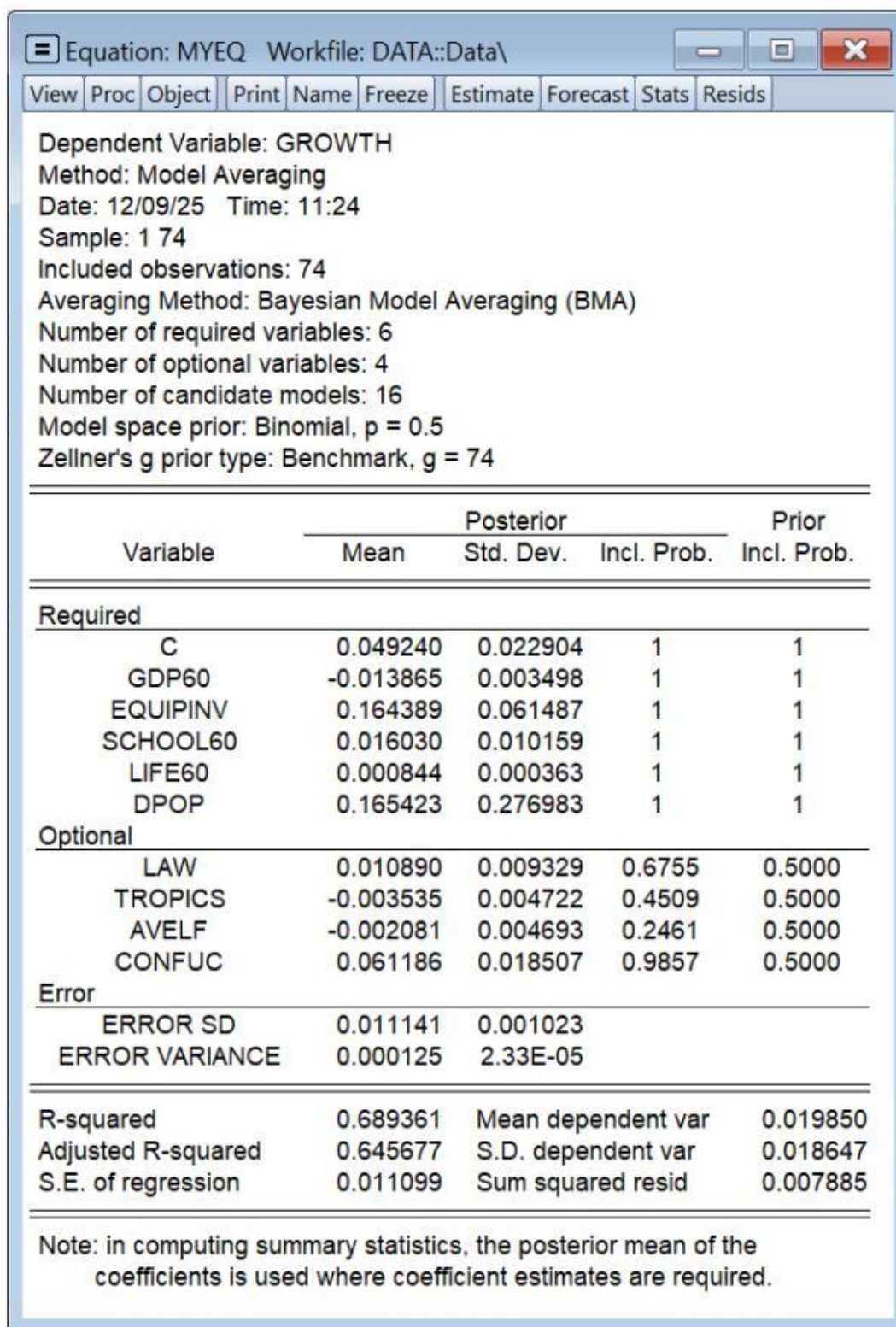
转到 **Quick/Estimate Equation...**，并从 **Method:** 下拉菜单中选择 **MODAVG - Model Averaging**。在 **Specification** 页面，在第一个编辑字段中输入

growth c gdp60 equipinv school60 life60 dpop

并在第二个字段中输入

law tropics avelf confuc

点击 **OK**，得到



“Posterior” 表头下的前两列显示了系数的后验均值与后验标准差。这些数值与 Magnus 等人 (2010) 表 2 最后一列所报告的值一致。

“Posterior” 表头下的第三列显示了后验包含概率 (PIP)。该列与 Magnus 等人 (2010) 表 6 正文中第二列一致。

新增或更新的命令

`modavg` 使用加权平均最小二乘 (WALS) 或贝叶斯模型平均 (BMA) 估计方程 (第 61 页)。(新增)

`modelranking` 显示按后验概率排序的模型 (第 62 页)。(新增)

`modelsizes` 显示与模型规模相关的结果 (第 63 页)。(新增)

`postmarginals` 显示参数后验分布的单变量边际 (第 68 页)。(新增)

弹性网络 GLM

EViews 15 在现有弹性网络估计的基础上进行了扩展，将广义线性模型 (GLM) 纳入 `enet` 工具集。

弹性网络 GLM 将正则化建模扩展到连续结果之外，通过恰当的分布与链接函数选择，使同一惩罚估计框架可应用于二值、计数、比例及其他非高斯响应变量。

Equation Estimation

Specification Penalty Options

Equation specification

Dependent variable followed by list of regressors:

Equation type: Least squares

Penalty specification

Type: Elastic net

Alpha: 0.9

Lambda:

(Enter a single lambda, or a list or vector of values to perform selection. Leave blank for EViews to supply a list.)

Estimation settings

Method: ENET - Elastic Net Regularization

Sample:

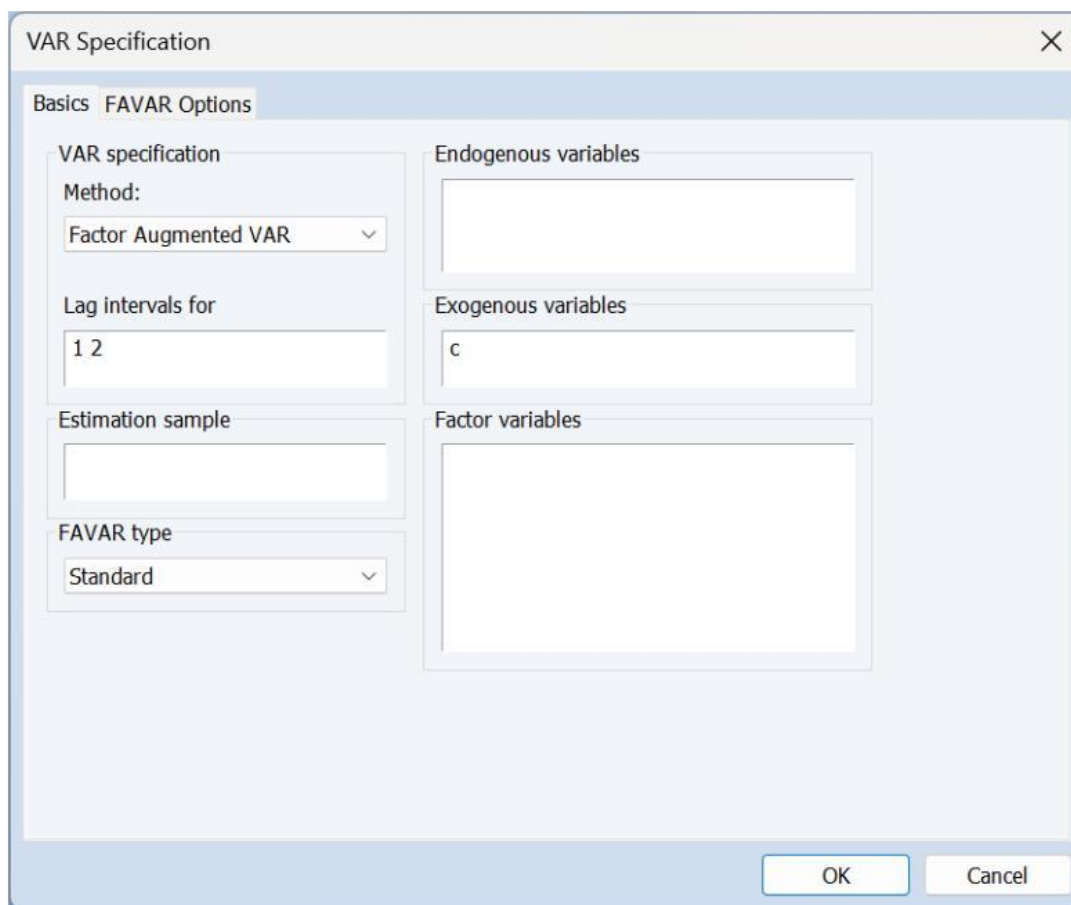
OK Cancel

新增或更新的命令

enet..... 弹性网络回归（含 Lasso 与岭回归）（第 39 页）。(已更新)

因子增广 VAR

因子增广 VAR（FAVAR）模型通过将一个由少量观测内生变量组成的系统与从更大一组相关序列中提取的信息相结合，扩展了标准 VAR 框架。FAVAR 并非将庞大的信息集直接纳入 VAR，而是使用少量估计得到的因子来概括因子变量的协同运动。随后，估计得到的因子会被追加到观测内生变量之后，并将该增广系统作为一个 VAR 来估计。



新增或更新的命令

`favar`.....估计因子增广 VAR 设定（第 50 页）。(新增)

BVAR 的增强

贝叶斯 VAR (BVAR) 将无约束 VAR 的似然函数与 VAR 参数的先验分布相结合。BVAR 建模之所以流行，是因为先验分布使建模者能够将拟合模型向某个基准模型收缩。在进行 VAR 分析时，收缩是可取的，有时甚至是必要的。

EViews 15 对 BVAR 做了以下更新：

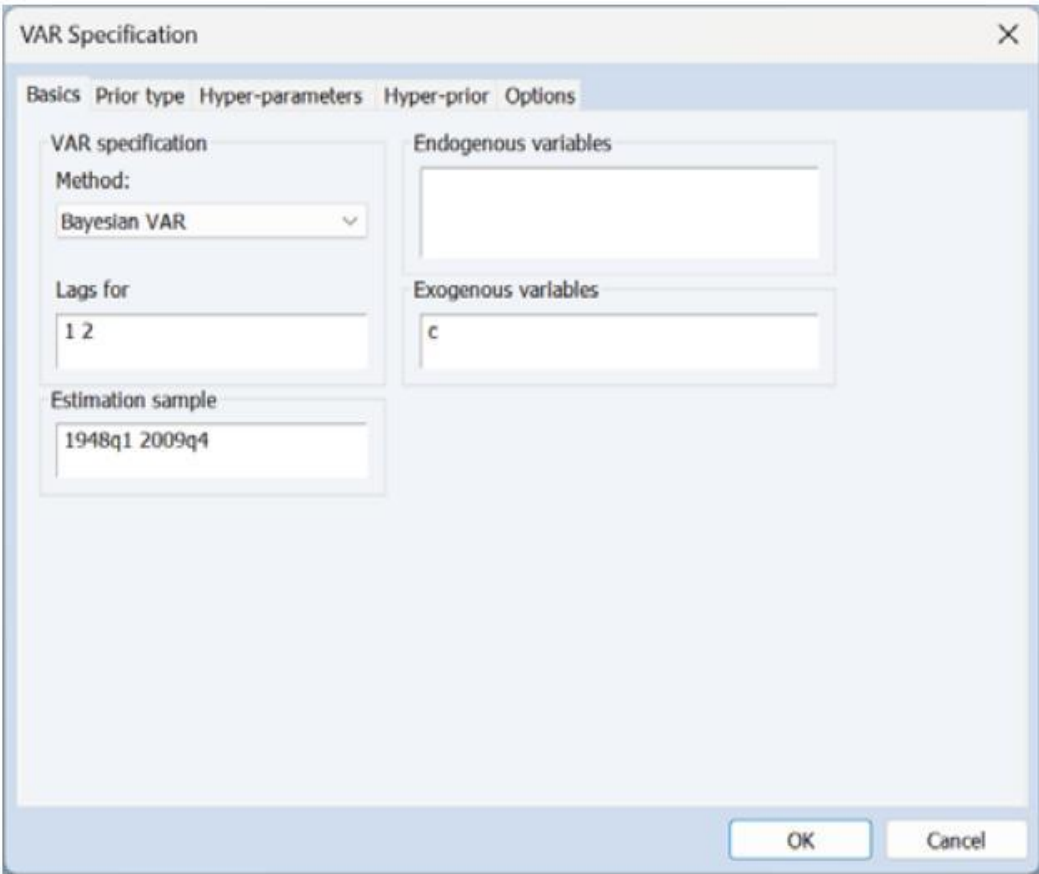
- 可用先验类型的选择已发生改变。
- 后验样本现已缓存，以便在估计后的流程中复用。
- GLP 已重新设计。

用户指南中关于 BVAR 的章节也经历了重大改动。有关 EViews 15 BVAR 更新的更多内容，请参见“版本兼容性说明”，起始于第 394 页。

(注意：早期版本的 EViews 无法读取在 EViews 15 中估计的 BVAR。反之，早期版本 EViews 中估计的 BVAR 需要在 EViews 15 中重新估计。)

要在 EViews 中估计 BVAR，请点击 **Quick/Estimate VAR...**，或在 command window 中输入 var 并按 Enter。VAR 对话框将打开。从 **Method:** 下拉菜单中选择 **Bayesian VAR**。VAR 对话框将更新以反映您的选择。

内生变量、外生变量、包含的滞后项以及估计样本在 **Basics** 页面中设置。



The screenshot shows the 'VAR Specification' dialog box with the 'Basics' tab active. The 'Method:' dropdown is set to 'Bayesian VAR'. The 'Lags for' field contains '1 2'. The 'Estimation sample' field contains '1948q1 2009q4'. The 'Endogenous variables' field is empty. The 'Exogenous variables' field contains 'c'. The 'OK' and 'Cancel' buttons are at the bottom right.

其余标签页允许您自定义 BVAR。

Prior type 页面允许您指定先验类型、是否包含虚拟观测 (dummy observations)，以及用于计算初始残差协方差矩阵 (如适用) 的设置。此外，若先验类型设置为 GLP，用户必须在此处选择需要推断的未知先验超参数。

The image shows a software window titled "VAR Specification" with a close button (X) in the top right corner. The window has five tabs: "Basics", "Prior type", "Hyper-parameters", "Hyper-prior", and "Options". The "Prior type" tab is currently selected. Inside this tab, there are three main sections:

- Prior type:** A dropdown menu set to "Litterman/Minnesota" and an unchecked checkbox labeled "L1 lag coefs only".
- Initial residual covariance options:** A section containing:
 - Model:** A dropdown menu set to "Univariate AR".
 - Exogenous variables:** A dropdown menu set to "Constant only".
 - Sample:** An empty text box with the instruction "(leave blank for estimation sample)" below it.
- Dummy observations:** A section containing:
 - Sum-of-coefficients:** A dropdown menu set to "Default".
 - Dummy initial observation:** A dropdown menu set to "Default".

At the bottom right of the window are "OK" and "Cancel" buttons.

先验超参数可在 **Hyper-parameters** 页面中设置。可用的先验超参数集合随其他设置（例如先验类型）而变化。

VAR Specification

×

Basics

Prior type

Hyper-parameters

Hyper-prior

Options

Prior base

Mu1:

1.0

AR(1) Coefficients

LambdaQ:

1.0

Residual covariance tightness*

Lambda1:

0.1

Overall tightness*

Lambda2:

0.99

Relative cross-variable weight

Lambda3:

1.0

Lag decay

Lambda4:

inf

Exogenous variables*

Q1:

0.1

Scaling factor on S

Q2:

0.1

Scaling factor on P

Q3:

Prior DOF (leave blank for default value)

Dummy observations

LambdaQ:

1.0

Sum-of-coefficients

LambdaZ:

1.0

Dummy initial observation

*You may use the keyword "inf" to specify infinite variance

OK

Cancel

当使用 GLP 时，超先验（即未知先验超参数上的分布）可在 **Hyper-prior** 页面中设置。

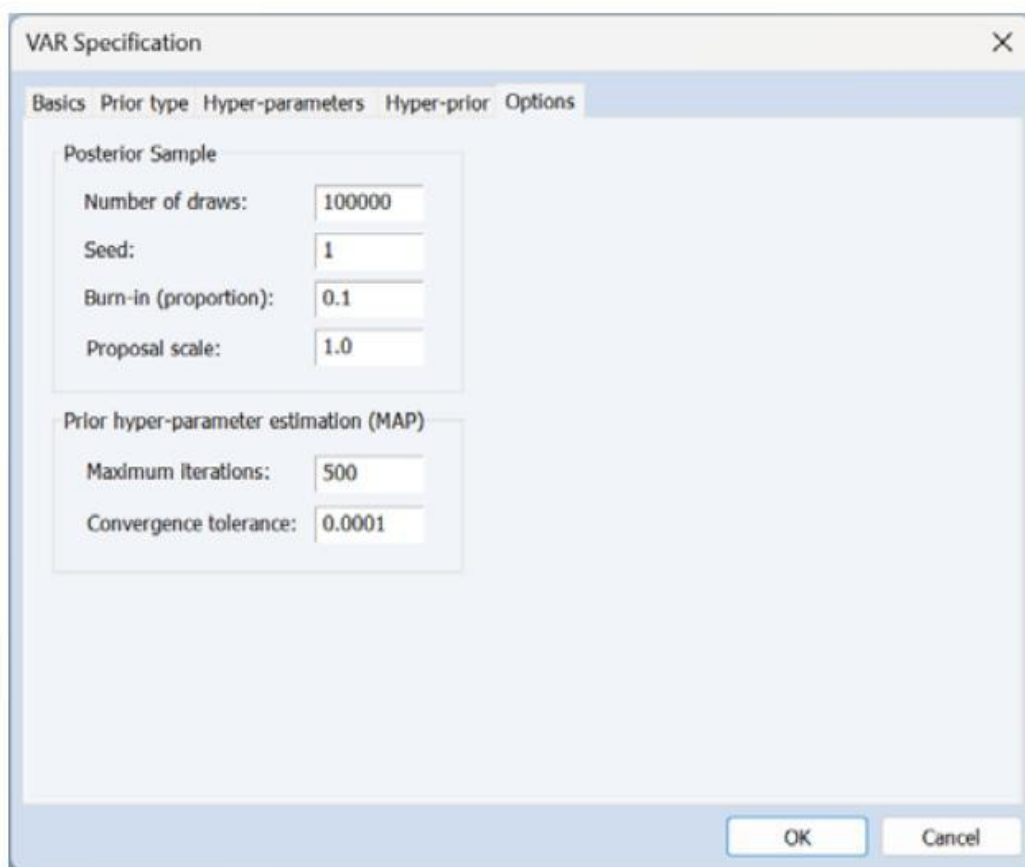
The image shows a software window titled "VAR Specification" with a close button (X) in the top right corner. The window has five tabs: "Basics", "Prior type", "Hyper-parameters", "Hyper-prior" (which is currently selected), and "Options".

Under the "Hyper-prior" tab, there are two sections:

- Initial prior**:
 - Lambda1:** Prior mode: 0.2, Prior Std. Dev.: 0.4
 - Psi:** Shape(s): 0.0004, Scale(s): 0.0004
- Dummy observations**:
 - Lambda6:** Prior mode: 1.0, Prior Std. Dev.: 1.0
 - Lambda7:** Prior mode: 1.0, Prior Std. Dev.: 1.0

At the bottom right of the window are two buttons: "OK" and "Cancel".

所有其余选项（例如与模拟相关的选项）在 **Options** 页面中设置。



更新后的 GLP 方法

GLP 方法已在 EViews 15 中重新设计。新增、修改或移除了若干选项，以使 EViews 对 GLP 的实现更好地与文献保持一致。

EViews 15 对 GLP 的更新包括：

- 能够指定超先验（即先验超参数上的分布）。
 - 对先验超参数采用不同的参数化。
 - 后验模拟的时机已移至估计期间进行。VAR 参数的估计值现变为后验均值，而非最大后验（MAP）估计。

GLP 方法示例

在本示例中，我们使用 `glpdata.wf1` 中的数据集，复现 Giannone 等人（2015）小规模 BVAR 的部分估计结果。

该小规模 BVAR 是作用于 GDP (Y)、GLP 平减指数 (P) 与联邦基金利率 (FFR) 的贝叶斯 VAR(5) 模型。为此，在 command window 中运行以下命令：

```
var bvar.bvar(prior=glp, l1glp, nsumcoef, ninitobs,  
    l4=@sqrt(10e6), l1lagcoefonly, initcov=ar1, icsmpl=1960q3  
    2008q4, propscale=50) 1 5 y p ffr
```

选项 prior=glp 与 l1glp 指示 EViews 使用 GLP 来推断 λ_1 。关键字 nsumcoef 与 ninitobs 用于排除虚拟观测；使用 GLP 时默认包含虚拟观测。选项 l4=@sqrt(10e6)、l1lagcoefonly、initcov=ar1、icsmpl=1960q3 2008q4 用于将我们的 BVAR 设置与论文保持一致。最后，选项 propscale=50 将提议分布 (proposal distribution) 的缩放因子设为 50，从而使 MH 接受率约为推荐的 20%。

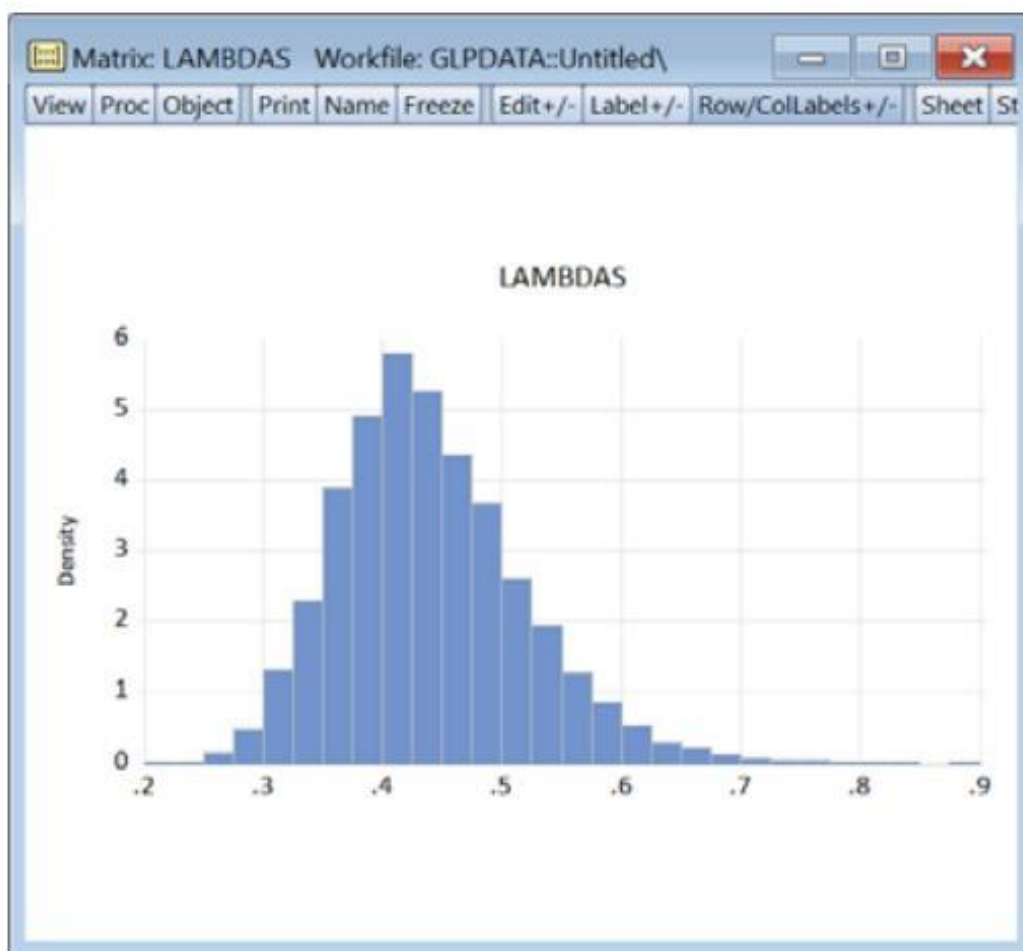
要显示 λ_1 的后验分布，先运行

```
bvar.postdraws coefs ecovs lambdas
```

将后验抽样保存到工作文件 (workfile)，然后运行

```
lambdas.distplot(s) hist(anchor=0, scale=dens)
```

所得直方图如下所示，对应于 Giannone 等人 (2015) 小规模 BVAR 图 1 中的曲线。



新增或更新的命令

bvar 估计贝叶斯 VAR 设定（第 32 页）。（已更新）
postdraws..... 将 BVAR 或 BTVCVAR 的后验抽样放入当前工作文件（第 6
7 页）。（已更新）
stablegr..... 执行稳定的 Gelman-Rubin (GR) 诊断（第 77 页）。
(新增)

增强的 Prophet

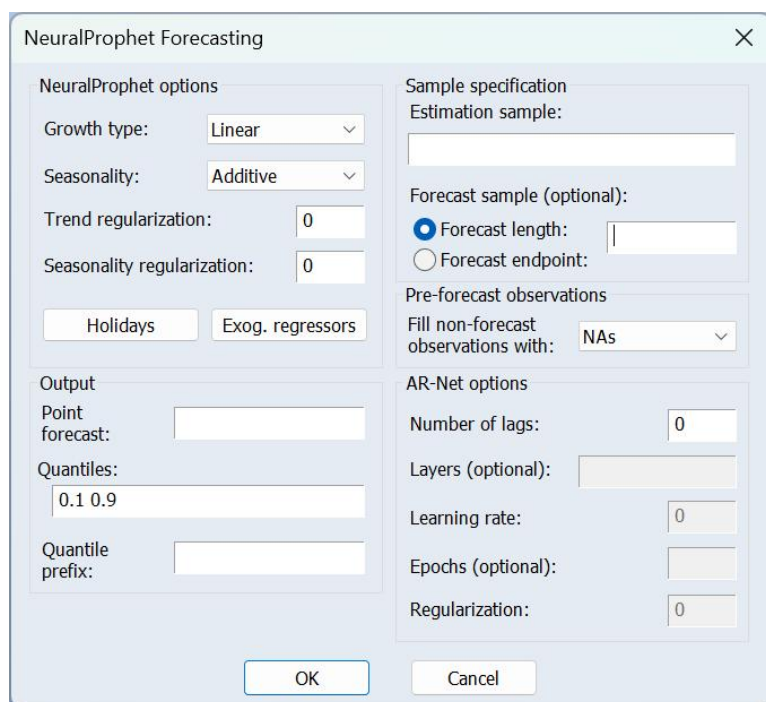
Prophet 已得到增强，允许指定节假日效应。

新增或更新的命令

prophet 对底层序列执行 Facebook 的 Prophet 预测（第 69 页）。(已更新)

NeuralProphet

NeuralProphet 是 Prophet 引擎的扩展，它将原始 Prophet 算法与更先进的深度学习技术相结合。具体而言，NeuralProphet 是一个混合模型，使用神经网络为 Prophet 增添自回归成分。



The image shows a 'NeuralProphet Forecasting' dialog box with the following sections:

- NeuralProphet options**
 - Growth type: Linear (dropdown)
 - Seasonality: Additive (dropdown)
 - Trend regularization: 0 (text input)
 - Seasonality regularization: 0 (text input)
 - Holidays (checkbox)
 - Exog. regressors (checkbox)
- Sample specification**
 - Estimation sample: (text input)
 - Forecast sample (optional):
 - ☒ Forecast length: (text input)
 - ☐ Forecast endpoint: (text input)
- Pre-forecast observations**
 - Fill non-forecast observations with: NAs (dropdown)
- Output**
 - Point forecast: (text input)
 - Quantiles: 0.1 0.9 (text input)
 - Quantile prefix: (text input)
- AR-Net options**
 - Number of lags: 0 (text input)
 - Layers (optional): (text input)
 - Learning rate: 0 (text input)
 - Epochs (optional): (text input)
 - Regularization: 0 (text input)

Buttons: OK, Cancel

新增或更新的命令

nprophet 对底层序列执行 NeuralProphet 预测（第 63 页）。(新增)

LSTM 模型

长短期记忆（LSTM）模型是一类为有序数据（sequential data）设计的循环神经网络，参见 Hochreiter 与 Schmidhuber（1997）、Gers、Schmidhuber 与 Cummins（2000）。在预测应用中，LSTM 学习从目标序列（以及可选的其他特征序列）的近期历史到目标序列未来值的非线性映射。EViews 提供了一个内置的

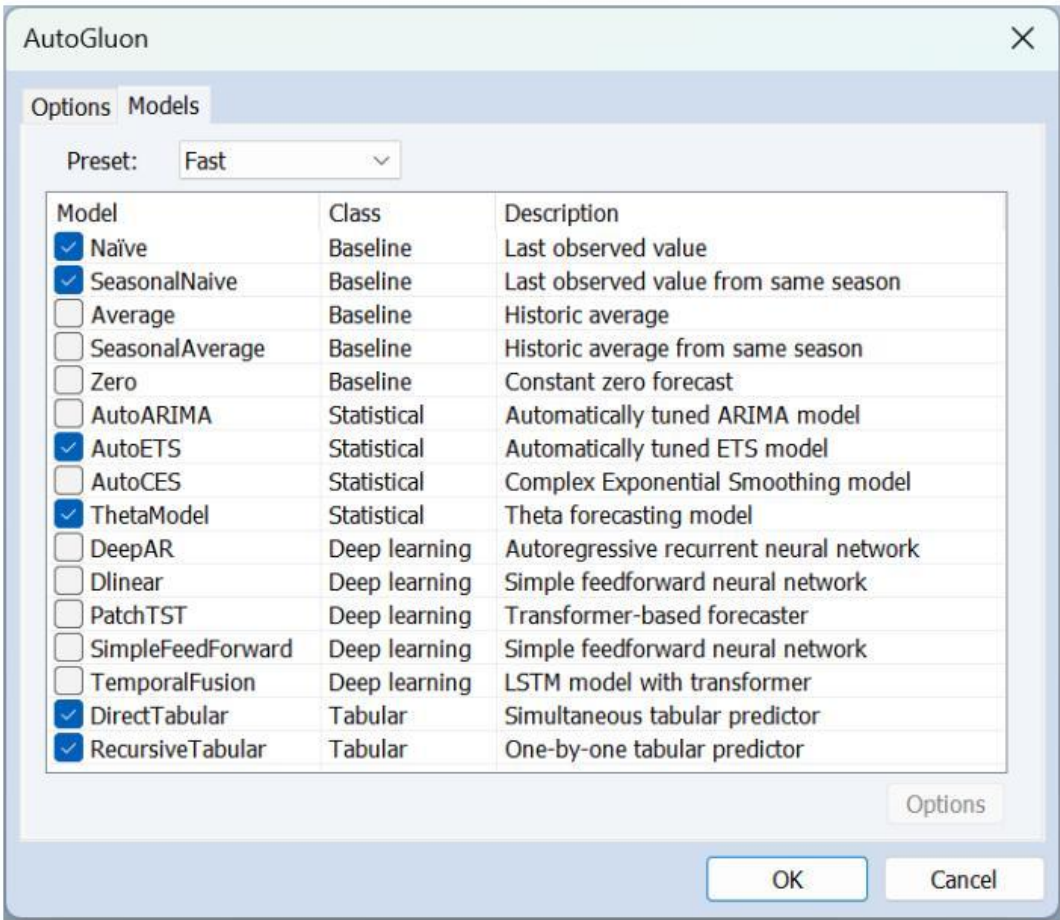
LSTM 预测引擎，可自动构建训练集与验证集，通过基于梯度的优化估计网络，并生成预测以及可选的诊断输出（损失图、拟合值与调参摘要）。

新增或更新的命令

lstm对底层序列执行长短期记忆估计与预测（第 54 页）。(新增)

AutoGluon

AutoGluon 是由 Amazon 设计的开源库，旨在简化机器学习模型的构建与优化过程。AutoGluon 最初专注于表格数据、图像分类与文本处理等任务，现已扩展其能力以支持时间序列预测，在一个软件包下提供了一套基于时间的预测模型。EViews 15 中可通过易用的界面访问 AutoGluon 提供的这套机器学习技术。



新增或更新的命令

autogluon..... 对底层序列执行 AutoGluon 预测（第 26 页）。(新增)

Lag-Llama

Lag-Llama 是构建于大语言模型（LLM）架构之上的自动预测软件包，旨在跨多种频率与预测期限提供准确的预测。其优势在于能够建模复杂动态，而这些动态对传统模型而言可能难以捕捉。

EViews 15 引入了易用的 Lag-Llama 界面，让用户能够便捷地使用其强大的预测方法。

The screenshot shows the 'Lag Llama' dialog box with the following settings:

- Options:**
 - Number of samples: 200
 - Context multiplier: 3
- Output:**
 - Point forecast: (empty)
 - Quantiles: 0.1 0.9
 - Quantile prefix: (empty)
- Sample specification:**
 - Estimation sample: (empty)
- Forecast:**
 - ☒ Forecast length: (empty)
 - ☐ Forecast endpoint: (empty)
- Fill pre-forecast observations with:** NAs

新增或更新的命令

lagllama.....对底层序列执行 Lag-Llama 预测（第 52 页）。(新增)

随机森林与决策树

随机森林（Random Forest）预测利用一组决策树的集成，通过对多个基于树的模型输出取平均来生成预测。随机森林对过拟合具有鲁棒性，调参直观，并能提供内

置的变量重要性度量，因而在传统线性方法不够用时，成为时间序列预测的一个颇具吸引力的选择。

Random forests and decision trees

Exogenous variables / factors

Lags: 0

☒ Drop variables with internal NAs

Sample specification

Estimation sample:

Forecast sample (optional):

☒ Forecast length:

☐ Forecast endpoint:

Options

Model: Random forest

Trees: 20

Leaf size: 10

Tree depth: 0

Max gain: 1e-07

Use bootstrap ☒

Random seed: 1

Learning rate: 0.1

Subsample: 1

Output

Series name:

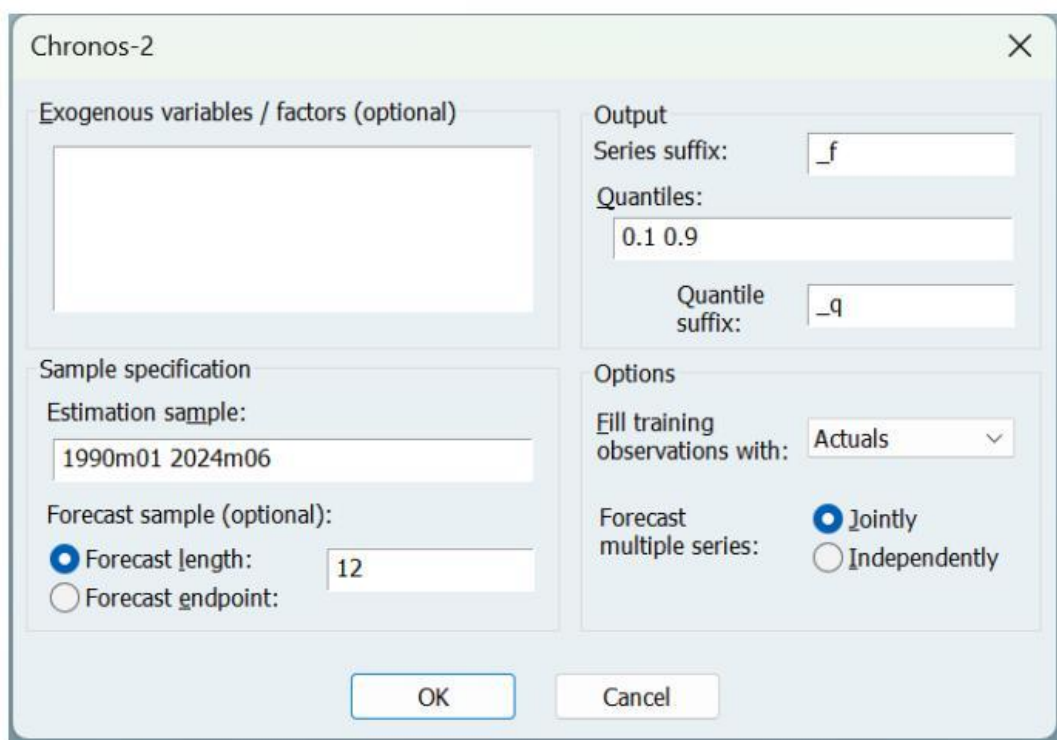
OK Cancel

新增或更新的命令

randomforest..... 执行随机森林风格的回归预测（第 71 页）。(新增)

Chronos-2

Chronos-2 是由 Amazon 开发的、基于 transformer 的开源基础模型，专为单变量与多变量时间序列预测而设计。与 Prophet、NeuralProphet 或 Lag-Llama 等其他自动预测模型不同，Chronos-2 采用先进的自然语言处理技术，将时间序列数据分词（tokenizing）为离散区间，并将未来值的预测作为一个语言建模任务。这种方法实现了零样本（zero-shot）预测能力，意味着它无需针对特定序列的训练即可生成准确的预测。



新增或更新的命令

`chronos2` 对底层序列执行 Chronos-2 预测（第 37 页）。（新增）
`chronos2` 对组（group）中的序列执行 Chronos-2 预测（第 36 页）。（新增）

增强的 ETS 平滑

ETS 平滑现已支持通过自助法（bootstrap）模拟生成预测区间。

新增或更新的命令

`ets`..... 执行误差-趋势-季节（ETS）估计与指数平滑（第 47 页）。（已更新）

第 4 章 命令参考的初步更新

本章包含 EViews 15 中新增或已更新命令的初步文档。

请注意，本文档为初步文档，且同样摘自一份更大的文档，因此部分内容可能格式不规范，指向页面与章节的交叉引用链接可能无法正常工作。

EViews 新增与更新命令的初步列表

方程

方程方法

enet..... 弹性网络回归（含 Lasso 与岭回归）（第 39 页）。(已更新)
modavg..... 使用加权平均最小二乘（WALS）或贝叶斯模型平均（BMA）估计方程（第 61 页）。(新增)
rdit..... 估计时间回归断点模型。（第 72 页）。(新增)
splinereg..... 使用非参数样条估计方程（第 76 页）。(新增)

方程视图

balance..... 查看使用协变量估计的时间回归断点模型的协变量平衡图（第 32 页）。(新增)
bwsensitivity..... 查看时间回归断点模型的带宽敏感性图（第 35 页）。(新增)
discontinuity..... 查看时间回归断点模型的断点图（第 38 页）。(新增)
donutsensitivity..... 查看时间回归断点模型的甜甜圈敏感性图（第 38 页）。(新增)
knots..... 查看使用样条估计的方程中各节点值的表（第 61 页）。(新增)
modelranking..... 显示按后验概率排序的模型（第 62 页）。(新增)
modelsize..... 显示与模型规模相关的结果（第 63 页）。(新增)
partialevects..... 查看使用样条估计的方程中偏效应的图形（第 66 页）。(新增)
placebo..... 查看时间回归断点模型的安慰剂图（第 66 页）。(新增)
postmarginals..... 显示参数后验分布的单变量边际（第 68 页）。(新增)

组

组过程 (Group Procs)

chronos2 对组中的序列执行 Chronos-2 预测 (第 36 页)。(新增)

序列 (Series)

序列过程 (Series Procs)

autogluon 对底层序列执行 AutoGluon 预测 (第 26 页)。(新增)

chronos2 对底层序列执行 Chronos-2 预测 (第 37 页)。(新增)

ets 执行误差-趋势-季节 (ETS) 估计与指数平滑 (第 47 页)。(已更新)

lagllama 对底层序列执行 Lag-Llama 预测 (第 52 页)。(新增)

lstm 对底层序列执行长短期记忆 (LSTM) 估计与预测 (第 54 页)。(新增)

nprophet 对底层序列执行 NeuralProphet 预测 (第 63 页)。(新增)

prophet 对底层序列执行 Facebook 的 Prophet 预测 (第 69 页)。(已更新)

splinedecom randomforest 执行随机森林风格的回归预测 (第 71 页)。(新增)使用样条分解对序列进行季节调整 (第 74 页)。(新增)

Var

Var 方法 (Var Methods)

bvarfavar 估计贝叶斯 VAR 设定 (第 32 页)。(已更新)估计因子增广 VAR 设定 (第 50 页)。(新增)

Var 过程 (Var Procs)

postdraws 将 BVAR 或 BTVCVAR 的后验抽样放入当前工作文件 (第 67 页)。(已更新)

Var 视图 (Var Views)

stablegr 执行稳定的 Gelman-Rubin (GR) 诊断 (第 77 页)。(新增)

autogluon 序列过程 (Series Procs)

对底层序列执行 AutoGluon 预测。

Prophet 需要随 EViews 安装的 AutoGluon Python 库。

语法

```
series_name.autogluon(options) forecast_name [quantile_prefix] [@modelspecs]
```

为预测序列输入一个名称，并可选择为分位数预测输入一个命名前缀，其后跟随 AutoGluon 所用各单个模型的设定说明。

可以使用 `preset=` 选项指定一组预置模型。如果包含了 `@modelspec` 参数，则该模型将与 `preset` 中指定的任意模型一同被纳入。若未使用 `preset` 选项，则仅使用由 `@modelspec` 参数指定的模型。

选项

- preset=arg** 选择一组预置模型。arg 可为 “naive”、“fast” (默认)、“medium”、“high”。
- eval=arg** 设定预测评估指标。arg 可为 SQL (缩放分位数损失, Scaled quantile loss)、WQL (加权分位数损失, Weighted quantile loss, 默认)、MAE (平均绝对误差, Mean absolute error)、MAPE (平均绝对百分比误差, Mean absolute percentage error)、MASE (平均绝对缩放误差, Mean absolute scaled error)、MSE (均方误差, Mean squared error)、RMSE (均方根误差, Root mean squared error)、RMSLE (均方根对数误差, Root mean squared logarithmic error)、SMAPE (对称平均绝对百分比误差, Symmetric mean absolute percentage error)、WAPE (加权绝对百分比误差, Weighted absolute percentage error)。
- quantiles=“arg”** 指定预测分位数。默认预测 0.1 与 0.9 分位数。
- f = arg**

预测样本选项

预测样本将从紧随估计样本（即当前工作文件样本）之后的第一个观测开始。预测终点由以下二者之一给出：

`forcen=int` 要预测的期数。

`forc="date"` 指定预测终点的日期。

若省略，则终点为工作文件范围（workfile range）的末尾。

模型设定

通过在 `autogluon` 命令之后以 `@` 符号包含各模型的关键字，可将单个模型添加到 AutoGluon 所估计的模型列表中。若某模型提供一组选项，可在关键字之后的括号中指定这些选项。模型关键字如下：

`@naive`

`@seasonalnaive`

`@average`

`@seasaverage`

`@zero`

`@autoarima(options)`

选项：

`diff=int` 设定差分的阶数。int 可取 0、1 或 2。若未指定，AutoGluon 将自动确定最合适的差分阶数。

`sdiff=int` 设定季节差分的阶数。int 可取 0 或 1。若未指定，AutoGluon 将自动确定最合适的季节差分阶数。

`maxar=int` 设定 AR 项的最大数量。默认值为 5。

`maxma=int` 设定 MA 项的最大数量。默认值为 5。

`maxsar=int` 设定季节 AR 项的最大数量。默认值为 0。

`maxsma=int` 设定季节 MA 项的最大数量。默认值为 0。

`nomean` 不包含截距项。

`stationary` 在搜索过程中仅纳入平稳 ARMA 模型。

`drift` 包含漂移项。

@*autoets*(options)

选项:

damped 对 ETS 模型的趋势分量进行阻尼。

@*autoces*

@*theta*(options)

选项:

decomp=arg 设定分解方式，可为乘法（mult，默认）或加法（add）。

@*deepar*(options)

选项:

context=int	设定循环神经网络展开的时间步数。
maxepoch=int	设定模型训练的 epoch 数量。
batchperepoch=int	设定每个 epoch 处理的训练数据批次数。
lr=num	设定训练过程的学习率/步长。
rnnlayers=int	设定循环神经网络中的层数。
rnncells=int	设定每层的单元数。
dropout=num	设定循环层训练期间使用的正则化参数。
noscale	不对每个上下文窗口应用均值绝对缩放。

@*dlinear*(options)

选项:

context=int	设定用于条件化预测的观测数量。
maxepoch=int	设定模型训练的 epoch 数量。
batchperepoch=int	设定每个 epoch 处理的训练数据批次数。
lr=num	设定训练过程的学习率/步长。
hiddenlayers=int	设定前馈网络中隐藏层的大小。
wgtdecay=num	设定权重衰减正则化参数。
scale=arg	设定训练与预测期间应用于每个上下文窗口的缩放方式。 arg 可为 “mean”（默认）、“stdev” 或 “none”。

@*patchtst*(options)

选项:

context=int	设定用于条件化预测的观测数量。
maxepoch=int	设定模型训练的 epoch 数量。
batchperepoch=int	设定每个 epoch 处理的训练数据批次数。
lr=num	设定训练过程的学习率/步长。
patchlen=int	设定每个补丁的长度。
stride=int	设定每个补丁的步幅。
hiddenlayers=int	设定 transformer 编码器中隐藏层的大小。
heads=int	设定 transformer 编码器中注意力头的数量。
wgtdecay=num	设定权重衰减正则化参数。

@simpfeedfwd(options)

选项:

context=int	设定用于条件化预测的观测数量。
maxepoch=int	设定模型训练的 epoch 数量。
batchperepoch=int	设定每个 epoch 处理的训练数据批次数。
lr=num	设定训练过程的学习率/步长。
hiddenlayers=int	设定前馈网络中隐藏层的大小。
batchnorm	在训练期间对每批输入进行归一化/标准化。
noscale	不对每个上下文窗口应用均值绝对缩放。

@tempfusion(options)

选项:

context=int	设定用于条件化预测的观测数量。
maxepoch=int	设定模型训练的 epoch 数量。
batchperepoch=int	设定每个 epoch 处理的训练数据批次数。
lr=num	设定训练过程的学习率/步长。
hiddenstates=int	设定 LSTM 与 transformer 隐藏状态的大小。
embedsize=int	设定嵌入的大小。

heads=int	设定 transformer 编码器中注意力头的数量。
-----------	-----------------------------

heads=int 设定 transformer 编码器中注意力头的数量。
dropout=num 设定循环层训练期间使用的正则化参数。
patience=int 设定早停限制。

@directtab(options)

选项:

scale=arg 设定应用于底层序列的缩放方式。arg 可为
 “mean”（默认）、“standard”或“none”。

@recurstab(options)

选项:

scale=arg 设定应用于底层序列的缩放方式。arg 可为
 “mean”（默认）、“standard”或“none”。

示例

```
elecddm.autogluon elecddm_f elec_f_q
```

这将在 ELECDMD 序列上执行 AutoGluon 例程，采用全部默认设置，将预测值放入序列 ELECDMD_F，并将 10% 与 90% 分位数分别存入 ELEC_F_Q_1 与 ELEC_F_Q_9。由于未指定 preset，因此使用默认的 Fast 预置，即采用 Naïve、Seasonal Naïve、Auto ETS、Theta、Direct Tabular 与 Recursive Tabular 方法。

```
elecddm.autogluon(preset=medium, f=actual) elecddm_f elec_f_q
```

指示 AutoGluon 使用 Medium 预置，即采用 Naïve、Seasonal Naïve、Auto ETS、Theta、Temporal Fusion、Direct Tabular 与 Recursive Tabular 方法。预测样本之外的观测以实际值而非 NA 插入预测序列。

```
elecddm.autogluon(preset=medium) elecddm_f elec_f_q  
@autoarima(maxma=3)
```

指示 AutoGluon 再次使用 Medium 预置，同时加入 Auto ARIMA 方法，并将最大 MA 项设为 3。

交叉引用

参见 *User's Guide I* 第 123 页的“AutoGluon”以获取更多讨论。

balance 方程视图

查看使用协变量估计的时间回归断点模型的协变量平衡图。

语法

eq_name.**balance**

示例

```
equation eq1.rdit log(conceptions) adoptions @ 2007m7  
eq1.balance
```

估计一个时间回归断点模型，以 CONCEPTIONS 的对数为结果变量、以 2007 年 7 月作为干预日期、以 ADOPTIONS 作为附加协变量，然后显示协变量平衡图。

交叉引用

参见 *User's Guide II* 第 273 页的“Regression Discontinuity in Time (RDiT)”以获取更多讨论。

bvar Var Methods

估计一个贝叶斯 VAR 设定。

语法

var_name.bvar(options) lag_pairs endog_list [@ exog_list]

bvar 估计一个贝叶斯 VAR。您必须指定 VAR 的阶数（使用一对或多对滞后区间），然后提供一个序列或组的列表作为内生变量。您可以通过包含“@”符号后接序列或组列表的方式，将趋势、季节虚拟变量等外生变量纳入 VAR。常数会自动添加到外生回归变量列表中；若要估计不含常数的设定，应使用“noconst”选项。

选项

通用选项

noconst	不在外生回归变量列表中包含常数。
prior = keyword(default = “lit”)	设定先验类型：“lit”（Litterman/Minnesota）、“nw”（Normal-Wishart）、“inw”（independent normal-Wishart）、“sz”（Sims-Zha）与“glp”（Giannone、Lenze 和 Primiceri）。若未指定先验类型，则使用

	Litterman/Minnesota 先验。
generic	使用先验的通用版本（仅适用于 normal-Wishart 与 independent normal-Wishart 先验）。若选项列表中不含关键字 “generic”，则使用 Litterman 型先验版本。
initcov = keyword(default = “uni”)	计算初始残差协方差矩阵所用的模型：“ar1”（单变量 AR(1)）、“uni”（单变量 AR）、“diag”（完整 VAR 的对角部分）与 “full”（完整 VAR）。若未指定，EViews 使用 “initcov=uni” 选项。
initexog = keyword(default = “const”)	设定用于计算初始残差协方差矩阵的外生回归变量列表：“none”（无外生回归变量）、“const”（仅常数）与 “all”（BVAR 中使用的全部外生回归变量）。若未指定，EViews 使用 “initexog=const” 选项。
icsmpl = arg	设定用于计算初始残差协方差矩阵的样本。若未指定样本，则使用估计样本。
sumcoef, nsumcoef	使用 “sumcoef” 纳入、“nsumcoef” 排除系数和（sum-of-coefficients）虚拟观测。若选项列表中二者均未包含，则使用先验特定的默认值。
initobs, ninitobs	使用 “initobs” 纳入、“ninitobs” 排除虚拟初始观测。若选项列表中二者均未包含，则使用先验特定的默认值。
l0 = num	设定误差协方差紧度（仅适用于 Sims-Zha 先验）。
l1 = num	设定总体紧度。
l1lagcoefonly	将总体紧度超参数（l1）仅应用于滞后系数。
l2 = num	设定相对交叉变量权重。
l3 = num	设定滞后衰减率。
l4 = num	设定外生变量紧度。
l6 = num	设定系数和（sum-of-coefficient）虚拟观测的尺度。
l7 = num	设定虚拟初始观测的尺度。
c1 = num	设定误差协方差矩阵的先验尺度。根据先验类型，可为标量、n 维向量或 n×n 对称矩阵。若使用标量或向量，则非对角元素设为零。

c1 = num	设定误差协方差矩阵的先验尺度。根据先验类型，可为标量、n 维向量或 $n \times n$ 对称矩阵。若使用标量或向量，则非对角元素设为零。
c2 = num	设定系数的先验尺度。根据先验类型，可为标量、k 维向量或 $k \times k$ 对称矩阵。若使用标量或向量，则非对角元素设为零。
c3 = num	设定误差协方差矩阵的先验自由度。默认值为
probstable	计算 VAR 平稳的后验概率。注意，这涉及使用后验样本。
l1glp	估计 l1 先验超参数。
l1mo=num	设定 l1 的超先验众数 (mode)。
l1sd=num	设定 l1 的超先验标准差。
l6glp	估计 l6 先验超参数。
l6mo=num	设定 l6 的超先验众数。
l6sd=num	设定 l6 的超先验标准差。
l7glp	估计 l7 先验超参数。
l7mo=num	设定 l7 的超先验众数。
l7sd=num	设定 l7 的超先验标准差。
psiglp	估计误差方差的先验均值。
psishapes=num	设定误差方差先验均值的超先验形状参数。
psiscales=num	设定误差方差先验均值的超先验尺度参数。
draws = num	设定要抽取的抽样次数。
seed = num	设定后验模拟器的随机种子。
burn = num	设定用于预烧 (burn-in) 的丢弃抽取百分比。
propscale=num	设定 Metropolis-Hastings (MH) 提议生成密度的尺度。仅适用于 GLP。
m=integer	设定优化的最大迭代次数。仅适用于 GLP。
c=scalar	设定优化的收敛准则。仅适用于 GLP。
prompt	强制在程序中弹出对话框。
p	打印基本估计结果。

示例

```
var var01.bvar 1 3 m1 gdp
```

声明并估计一个名为 VAR01 的 VAR 对象，它是一个使用 Litterman/Minnesota 先验的贝叶斯 VAR，具有两个内生变量（M1 与 GDP）、一个常数以及 3 阶滞后（1 至 3）。

```
var01.bvar(noconst) 1 3 m1 gdp
```

估计相同的 BVAR，但省略常数。

```
var var02.bvar(prior=nw, m1=0.5) 1 3 m1 gdp
```

声明并估计一个使用（共轭）normal-Wishart 先验的 BVAR，其先验超参数 m1 设为 0.5。

```
var var03.bvar(prior=inw, m1=0.5) 1 3 m1 gdp
```

在 independent normal-Wishart 先验下执行相同的操作。

```
var var04.bvar(prior=glp, l1glp) 1 3 m1 gdp
```

使用 GLP 先验，不仅估计 VAR 参数，还估计总体紧度超参数 l1。

交叉引用

详见 *User's Guide II* 第 369 页第 12 章“Bayesian VAR Models”。

bwsensitivity Equation Views

查看时间回归断点模型的带宽敏感性图。

语法

```
eq_name.bwsensitivity
```

示例

```
equation eq1.rdit log(conceptions) @ 2007m7  
eq1.bwsensitivity
```

估计一个时间回归断点模型，以 CONCEPTIONS 的对数为结果变量、以 2007 年 7 月作为干预日期，然后显示带宽敏感性图。

交叉引用

参见 *User's Guide II* 第 273 页的“Regression Discontinuity in Time (RDiT)”以获取更多讨论。

chronos2 Group Procs

对组（group）中的序列执行 Chronos-2 预测。

本命令需要安装随 EViews 提供的 Chronos-2 Python 库。

语法

`group_name.chronos2(options) forecast_suffix [quantile_suffix] @ exog`

输入预测序列的命名模式后缀，并可选择输入分位数预测的命名模式后缀。若希望纳入外生变量，请使用“@”后接序列名称列表。

选项

<code>independent</code>	对序列执行独立的单变量预测。若省略，Chronos-2 将执行联合预测。
<code>quantiles="arg"</code>	指定预测分位数。默认预测 0.1 与 0.9 分位数。
<code>f=arg</code>	预测样本外填充行为：“actual”（以拟合变量的实际值填充预测样本之外的观测），或“na”（以缺失值填充预测样本之外的观测，默认）。

预测样本选项

预测样本将从估计样本（即当前工作文件样本）之后的第一个观测开始。预测终点由以下二者之一给出：

<code>forclen=int</code>	要预测的期数。
<code>forc="date"</code>	指定预测终点的日期。

若省略，则终点为工作文件范围的末尾。

示例

```
macrovars.chronos2 _f _f_q @ @trend
```

将对组（group）MACROVARS 中的序列执行 Chronos-2 预测，将预测值放入与原始序列同名但追加 _F 的序列中，并将 10% 与 90% 分位数分别存入 *_F_Q_1 与 *_F_Q_9。时间趋势作为外生变量纳入。

交叉引用

参见 *User's Guide I* 第 170 页的“Chronos-2”以获取更多讨论。

chronos2 Series Procs

对底层序列执行 Chronos-2 预测。

本命令需要安装随 EViews 提供的 Chronos-2 Python 库。

语法

```
series_name.chronos2(options) forecast_name [quantile_prefix] @ e  
xog
```

为预测序列输入一个名称，并可选择为分位数预测输入一个命名前缀。若希望纳入外生变量，请使用“@”后接序列名称列表。

选项

quantiles="arg"	指定预测分位数。默认预测 0.1 与 0.9 分位数。
f=arg	预测样本外填充行为：“actual”（以拟合变量的实际值填充预测样本之外的观测），或“na”（以缺失值填充预测样本之外的观测，默认）。

预测样本选项

预测样本将从估计样本（即当前工作文件样本）之后的第一个观测开始。预测终点由以下二者之一给出：

forclen=int 要预测的期数。
forc="date" 指定预测终点的日期。

若省略，则终点为工作文件范围的末尾。

示例

```
elec dmd.chronos2 elec dmd_f elec_f_q @ temp @expand(@month)
```

这将在 ELECDMD 序列上执行 Chronos-2 例程，将预测值放入序列 ELECDMD_F，并将 10% 与 90% 分位数分别存入 ELEC_F_Q_1 与 ELEC_F_Q_9。序列 TEMP 以及月度虚拟变量用作外生变量。

交叉引用

参见 *User's Guide I* 第 170 页的“Chronos-2”以获取更多讨论。

discontinuity Equation Views

查看时间回归断点模型的断点图。

语法

```
eq_name.discontinuity
```

示例

```
equation eq1.rdit log(conceptions) @ 2007m7  
eq1.discontinuity
```

估计一个时间回归断点模型，以 CONCEPTIONS 的对数为结果变量、以 2007 年 7 月作为干预日期，然后显示断点图。

交叉引用

参见 *User's Guide II* 第 273 页的“Regression Discontinuity in Time (RDiT)”以获取更多讨论。

donutsensitivityEquation Views

查看时间回归断点模型的甜甜圈敏感性图。

语法

```
eq_name.donutsensitivity
```

示例

```
equation eq1.rdit log(conceptions) @ 2007m7  
eq1.donutsensitivity
```

估计一个时间回归断点模型，以 CONCEPTIONS 的对数为结果变量、以 2007 年 7 月作为干预日期，然后显示甜甜圈敏感性图。

交叉引用

参见 *User's Guide II* 第 273 页的“Regression Discontinuity in Time (RDiT)”以获取更多讨论。

enet Equation Methods

弹性网络模型的估计，包含用于 Lasso 与岭回归（ridge regression）的选项。

语法

```
equation_name.enet(options) y x1 [x2 x3 ...] [@vw(...)]
```

先列出因变量，后接自变量列表。若希望包含不受罚的截距项，请使用“C”。

注意，弹性网络设定中不允许 PDL 与 ARMA 项。

如果您希望为回归变量指定单独的罚权重，或对系数值施加不等式约束，可以使用如下形式的特殊表达式（其中记变量 j 的罚权重为 q_j ）：

```
@vw(series_name, weight_value)
```

或

```
@vw(series_name[, wgt=weight_value, cmin=coef_min, cmax=coef_max])
```

其中 *weight_value* 为非负值，* *coef_min* * 为非正的最小系数值，* *coef_max* * 为非负的最大系数值。

该特殊表达式有两种形式。

在简写形式中，您指定变量名称后接罚权重值。

在更一般的形式中，您指定变量名称后接一个或多个按任意顺序排列的关键字表达式，其中“wgt=”参数指定罚权重，“cmin=”后接非正的最小系数值，“cmax=”后接非负的最大系数值。

使用 @vw 指定单个回归变量行为时，请注意：

- 特殊截距变量“C”始终不受罚，并隐含权重 $\omega = 0.0$ 。
- 单个罚权重也可使用 **Penalty** 对话框页面上的 **Individual lambda wgt.** **vector** 编辑字段（或使用命令估计选项 “`lambdawgt=vector_name`”）来指定。若同时指定了向量权重与使用 @vw 关键字指定的单个权重，则先应用向量权重，再应用单个变量权重。
- 单个系数限值也可使用 **Options** 对话框页面上的 **Min values vector** 与 **Max values vector** 编辑字段（或命令估计选项 “`coefmin=vector_name`” 与 “`coefmax=vector_name`”）以向量形式指定。若同时指定了向量系数限值与使用 @vw 关键字指定的单个回归变量限值，则先应用向量限值，再应用单个限值。

EViews 会将各单个罚权重归一化，使其之和等于系数个数。

选项

设定选项

eqtype=arg	估计广义线性模型（GLM）；若省略则为普通线性模型。
penalty=arg (default="el")	阈值估计的类型：“enet”（弹性网络，elastic net）、“ridge”（岭回归，ridge）、“lasso”（Lasso）。
alpha=arg (default=".5")	混合参数的值。必须为 0 到 1 之间的值。
lambda=arg	罚参数的值。可以是一个或多个数值或向量对象。值必须大于或等于零。若留空（默认），EViews 将生成一个列表。

GLM 选项

以下设定选项仅在方程类型设为广义线性模型（GLM）时适用。

family=arg (default="normal")	分布族：正态（“normal”）、泊松（“poisson”）、二项计数（“binomial”）、二项比例（“binprop”）、负二项（“negbin”）、伽马（“gamma”）、逆高斯（“igauss”）、指数均值（“emean”）、幂均值（“pmean”）、二项平方（“binsq”）。二项计数、二项比例、负二项与幂均值族均需要指定一个分布参数：
n=arg	二项计数（“family=binomial”）或二项比例

(default=1)	("family=binprop") 族的试验次数。
fparam = arg	负二项 ("family = negbin") 与幂均值 ("family = pmean") 族的分布族参数值。
link = arg (default = "identity")	链接函数：恒等 ("identity")、对数 ("log")、对数补集 ("logc")、Logit ("logit")、Probit ("probit")、对数-对数 ("loglog")、互补对数-对数 ("cloglog")、倒数 ("recip")、幂 ("power")、Box-Cox ("boxcox")、幂优势比 ("opow")、Box-Cox 优势比 ("obox")。幂、Box-Cox、幂优势比与 Box-Cox 优势比链接均需要指定一个链接参数，使用 "lparam =" 来设定。
lparam = arg	幂 ("link = power")、Box-Cox ("link = boxcox")、幂优势比 ("link = opow") 与 Box-Cox 优势比 ("link = obox") 链接函数的链接参数。
offset = arg	偏移项。
disp = arg	离散度估计量：Pearson χ^2 统计量 ("pearson")、偏差统计量 ("deviance")、单位 ("unit")、用户指定 ("user")。默认值依族而定：二项计数、二项比例、负二项与泊松为 "unit"，其余为 "pearson"。"deviance" 选项仅对指数分布族中的族提供（正态、泊松、二项计数、二项比例、负二项、伽马、逆高斯）。
dispval = arg	用户离散度值（若 "disp = user"）。
fwgts = arg	频率权重。
w = arg	权重序列或表达式。
wtype = arg (default = "istdev")	权重指定类型：逆标准差 ("istdev")、逆方差 ("ivar")、标准差 ("stdev")、方差 ("var")。
wscale = arg	权重缩放：EViews 默认 ("evIEWS")、平均 ("avg")、无 ("none")。默认设置取决于权重类型：若 "wtype = istdev" 则为 "evIEWS"，其余为 "avg"。

罚选项

ytrans = arg (default = "none")	因变量的缩放方式: "none" (无)、"L1" (L1)、"L2" (L2)、"stdsmpl" (样本标准差)、"stdpop" (总体标准差)、"minmax" (最小值-最大值)。
xtrans = arg (default = "stdpop")	回归变量的缩放方式: "none" (无)、"L1" (L1 范数)、"L2" (L2 范数)、"stdsmpl" (样本标准差)、"stdpop" (总体标准差)、"minmax" (最小值-最大值)。
nlambdas = integer (default = 100)	EViews 提供列表的罚值数量。
lambdaratio = arg	EViews 提供列表中最小与最大 lambda 之比。您可以为比例参数指定一个值, 也可以将编辑字段留空, 让 EViews 根据观测数 T 与潜在回归变量数 p 指定默认值。默认情况下, EViews 将该比例设为 0.001。
lambdawgt = vector_name	单个罚权重的向量, 包含非负值, 其大小与设定中的变量相匹配且顺序与之一致。若提供了向量且使用一个或多个 (vw?) 回归变量指定了单个权重, 则先应用向量权重, 再被单个变量权重覆盖。为便于比较, 我们将最终权重归一化, 使其之和为 p^* , 其中 p^* 为非零 ω_j 的个数。
nlambdamin = integer (default = 5)	应用停止规则前路径中的 lambda 值的最小数量。
minddev = arg (default = 1e-05)	继续估计所需的偏差分数最小变化。若偏差的相对变化小于此值, 则截断路径估计。
maxedev=arg (default=0.99)	终止估计时所达到的已解释偏差分数的最大值。若已解释的零模型偏差 (null deviance) 比例大于此值, 则截断路径估计。
maxvars=arg	模型中回归变量的最大数量。若系数数量 (含截距等非受罚变量的系数) 达到此值, 则截断路径估计。
maxvarsratio=arg	模型中回归变量的最大数量, 以观测数的比例表示。若系数数量 (含截距等非受罚变量的系数) 除以观测数达到此值, 则截断路径估计。

cvmethod=arg (default="kfold_cv")	交叉验证方法: "kfold" (k 折)、"simple" (简单划分)、"mcarlo" (蒙特卡洛, Monte Carlo)、"leavepout" (留 P 法, leave-P-out)、"leave1out" (留一法, leave-1-out)、"rolling" (滚动窗口)、"expanding" (扩展窗口)。
cvmeasure=arg (default="mse")	交叉验证拟合度量: "mse" (均方误差, mean-squared error)、"r2" (R-squared)、"mae" (平均绝对误差, mean absolute error)、"mape" (平均绝对百分比误差, mean absolute percentage error)、"smape" (对称平均绝对百分比误差, symmetric mean absolute percentage error)。
cvnfolds=arg (default=5)	K 折交叉验证的折数。 适用于 "cvmethod=kfold"。
cvftrain=arg (default=0.8)	简单划分与蒙特卡洛方法所用数据的比例。 适用于 "cvmethod=simple" 与 "cvmethod=mcarlo"。
cvnreps=arg (default=1)	蒙特卡洛方法的重复次数。 适用于 "cvmethod=mcarlo"。
cvleaveout=arg (default=2)	留 P 法所留出的数据点数量。 适用于 "cvmethod=leavepout"。
cvnwindows=arg (default=4)	滚动窗口交叉验证方法的窗口数量。适用于 "cvmethod=rolling"。
cvinitial=arg (default=12)	扩展交叉验证训练集中初始数据点数量。适用于 "cvmethod=expanding"。

交叉验证选项

cvpregap=arg (default=0)	训练集末尾与测试集开始之间的观测数量。适用于 "cvmethod=simple"、"cvmethod=rolling" 与 "cvmethod=expanding"。
cvhorizon=arg (default=1)	测试集中的观测数量。适用于 "cvmethod=rolling" 与 "cvmethod=expanding"。
cvpostgap=arg	滚动窗口下测试集末尾与下一训练集开始之间、或扩展窗口下

(default=0) 测试集末尾与下一训练集末尾之间的观测数量。适用于
"cvmethod=rolling" 与 "cvmethod=expanding"

随机数选项

seed=positive_integer
from 0 to 2,147,483,647 为随机数生成器设定种子。若未指定，EViews 将以默认全局随机数生成器的一次整数抽取结果为其设定种子。

rnd=arg
(default="kn" or method
previously set using
rndseed(p. 576) in the
Command and
Programming Reference). 随机数生成器类型：改进的 Knuth 生成器 ("kn")、改进的 Mersenne Twister ("mt")、EViews 4 中使用的 Knuth (1997) 滞后斐波那契生成器 ("kn4")、L'Ecuyer (1999) 组合多重递归生成器 ("le")、Matsumoto 与 Nishimura (1998) 的 Mersenne Twister ("mt4")。

其他选项

coefmin=vector_name,
number 单个系数最小值的向量，包含负数或缺失值，其大小与设定中的变量相匹配且顺序与之一致；或用于所有系数的最小值的负值。

coefmax=vector_name,
number 各系数最大值的向量，包含与设定中变量顺序一致且大小匹配的正值或缺失值，或为一个正值，表示所有系数的最大值。向量中的缺失值应用于表示对应的系数不受限制。若提供数值向量，且使用一个或多个 @vw 回归变量指定了各自的最大值，则先应用向量中的值，随后被各自的值覆盖。

maxit=integer 最大迭代次数。

conv=scalar 设定收敛准则。该准则基于缩放后估计值百分比变化的最大值。EViews 将该准则设为 1e-24 与 0.2 之间最接近的值。

w=arg 加权序列或表达式。

wtype=arg(default="istdev") 加权设定类型：逆标准差 ("istdev")、逆方差 ("ivar")、标准差 ("stdev")、方差 ("var")。

wscale=arg 权重缩放：EViews 默认 ("eviews")、平均 ("avg")、无 ("none")。默认设置取决于加权类型：若 "wtype=istdev" 则为 "eviews"，其余为 "avg"。

estmethod=arg	估计方法：“cov”（使用协方差算法）或“naive”（使用朴素算法）。默认为 EViews 自动选择。
showopts / -showopts	[显示 / 不显示] 输出中的估计选项。
prompt	在程序内部强制弹出对话框。
p	估计后打印基本结果视图。

示例

命令

```
equation enet1.enet(xtrans=none, lambdaratio=.0001,
    cvseed=513255899) lpsa c lcavol lweight age lbph svi lcp
gleason pgg45
```

估计一个弹性网络模型， α 取默认值 0.9，回归变量与因变量均不进行缩放，自动确定由 100 个元素组成的 lambda 路径，最小 lambda 为最大值的 0.0001 倍，使用默认的 5 折（K-fold）交叉验证，并以 MSE 为目标函数，使用随机数生成器种子 513255899 来确定最优值。

类似地，

```
equation lasso1.enet(penalty=lasso, lambdaratio=.0001,
    cvseed=513255899) lpsa c lcavol lweight age lbph svi lcp
gleason pgg45
```

估计一个 Lasso 模型，对回归变量进行总体标准差缩放，其余设定同上；而

```
equation ridge1.enet(penalty=ridge, lambdaratio=.0001,
    cvseed=513255899) lpsa c lcavol lweight age lbph svi lcp
gleason pgg45
```

估计等价的岭回归设定。

命令

```
equation enet2.enet(alpha=0.75, lambdaratio=.0001,
    cvmethod=rolling, cvmeasure=smape) lpsa c lcavol lweight age_s
lbph svi lcp gleason pgg45
```

估计一个弹性网络模型， α 等于 0.75，使用滚动窗口交叉验证与 SMAPE 交叉验证。

我们可以使用 (**vw?**) 设定来分配各自独立的罚项与系数限制：

```
equation enet3.enet(alpha=0.75, lambdaratio=.0001,
    cvmethod=rolling, cvseed=513255899) lpsa c @vw(lcavol, cmax=.4)
lweight age lbph svi lcp gleason @vw(pgg45, cmax=0.075, w=1.2)
```

估计一个弹性网络模型，其中 LCAVOL 的系数被限制为小于或等于 0.4，PGG45 的系数具有相对罚项权重 1.2，且最大值为 0.075。

等价的设定也可以使用罚项权重向量与系数限制向量来估计：

```
vector(9) cmax = na
cmax(2) = 0.4
cmax(9) = 1.2
vector(9) lwgt = 1
lwgt(9) = 1.2
equation enet4.enet(alpha=0.75, coefmax=cmax, lambdawgt=lwgt,
    lambdaratio=.0001, cvmethod=rolling, cvseed=513255899) lpsa c
    lcavol lweight age lbph svi lcp gleason pgg45
```

以及

```
vector(9) cmax = na
cmax(2) = 0.4
vector(9) lwgt = 1
lwgt(9) = 50
equation enet5.enet(alpha=0.75, coefmax=cmax, lambdawgt=lwgt,
    lambdaratio=.0001, cvmethod=rolling, cvseed=513255899) lpsa c
    lcavol lweight age lbph svi lcp gleason @vw(pgg45, cmax=0.075,
    w=1.2)
```

因为向量中 PGG45 的罚项权重会被使用@**VW** 指定的各自权重覆盖。

注意，无论哪种情形，截距均不受到惩罚，即使 LWGT 的对应元素等于 1，因为“C”的设定总是被隐式地当作“@**VW**(C, 0)”处理。

交叉引用

参见《User's Guide II》第 303 页第 10 章“Elastic Net and Lasso”，其中讨论了线性模型与广义线性模型的弹性网络、岭回归与 Lasso 估计。

etsSeries Procs

执行误差 - 趋势 - 季节 (ETS) 指数平滑

ets 过程采用 ETS 模型框架对序列进行预测，该框架基于状态空间的似然计算，支持模型选择，并可计算预测标准误。

ETS 框架定义了一类扩展的指数平滑模型，包括标准指数平滑模型（例如 Holt 与 Holt-Winters 的加法与乘法模型）。

语法

`series_name.ets(options) smooth_name`

应输入 `ets` 关键字，后接选项，再输入平滑输出序列的名称。您可以在括号中指定平滑方法（默认设置为加法误差、无趋势、无季节性）与平滑选项。

选项

预测样本选项

预测样本从估计样本（即当前工作文件样本）之后的第一个观测值开始。预测终点由以下二者之一给定：

`forclen=int` 预测期数。

`forc="date"` 指定预测终点的日期。

二者之一为必填项。

常规

`prompt` 在程序内部强制弹出对话框。

`p` 打印视图。

模型设定

`e=arg` 设定误差类型：“a”（加法）、“m”（乘法）、“e”（自动）。
(default = "a")

`t=arg` 设定趋势类型。键值可为：“n”（无）、“a”（加法）、“m”（乘法）、“ad”（加法阻尼）、“md”（乘法阻尼）、“e”（自动）。
(default = "n")

`s=arg` 设定季节类型。键值可为：“n”（无）、“a”（加法）、“m”（乘法）、“e”（自动）。
(default = "n")

`model=arg` 模型选择方法：“aic”（赤池信息准则，Akaike information criterion）、“bic”（贝叶斯信息准则/施瓦茨准则，Bayesian information criterion/Schwarz criterion）、“hq”（汉南-奎因信息准则，Hannan-Quinn information criterion）、“amse”（平均均方误差，average mean squared errors）。
(default= "aic")

`alpha=arg` 为水平参数 (α) 指定固定值。

beta=arg	为含趋势模型中的趋势参数 (β) 指定固定值。
gamma=arg	为含季节成分模型中的季节参数 (γ) 指定固定值。
phi=arg	为含阻尼趋势模型中的阻尼参数 (ϕ) 指定固定值。
nomult	不允许乘法趋势或季节项。仅在设置了 t=e 或 s=e 选项时适用。

优化选项

amse	将平均均方误差 (AMSE) 设为目标函数 (默认以对数似然为目标函数)。
namse=integer	指定 AMSE 长度——若选择 "amse", 用于计算 AMSE 的观测值个数。
c=number	设定收敛准则。
m=integer	设定最大迭代次数。
ustart	采用用户提供的初值 (取自工作文件中的 C 向量)。
noi	不对初始状态值进行优化 (固定为其初始值)。

输出选项

dgraph=arg	为每个指定元素包含分解图。arg 可由下列任意元素组成: "f" (预测, forecast)、"l" (水平, level)、"t" (趋势, trend)、"s" (季节, season)。
dgopt=arg(default="m")	分解图的显示格式: "m" (多图, multiple graph)、"s" (单图, single graph)。
graph=arg	在输出中为每个指定元素包含比较图 (若采用模型选择)。arg 可由下列元素组成: "c" (预测比较, forecast comparison) 与 "l" (似然比较, likelihood comparison)。
table=arg	在输出中包含比较表 (若采用模型选择)。arg 可由下列元素组成: "c" (预测比较) 与 "l" (似然比较)。
level=name	将水平分量作为独立序列保存到工作文件中。
trend=name	将趋势分量作为独立序列保存到工作文件中 (如适用)。
season=name	将季节分量作为独立序列保存到工作文件中 (如适用)。

自助法选项

<code>boot</code>	执行自助法 (Bootstrap) 重复抽样以得到预测区间。
<code>nsamples=int</code> (default=0)	自助法重复抽样次数。
<code>bootseed=int</code> (default=1)	自助法重抽样所用随机数生成器的种子。
<code>quantiles=arg</code> (default="0.1 0.9")	要计算的模拟预测分布分位数的空格分隔列表。每个值必须严格介于 0 与 1 之间。
<code>quantprefix=name</code>	用于命名所保存分位数序列的前缀。若未指定，则以 "Q" 作为前缀。

示例

```
sales.ets(e=a, t=n, s=a)sales_f
```

使用 ANN (加法误差、无趋势、无季节) 模型对序列 SALES 进行平滑，并创建名为 "SALES_F" 的平滑序列。

```
tb3.ets(e=e, t=e, s=n) tb3_smooth
```

将对 TB3 进行平滑，在不同误差与趋势设定之间自动选择最优平滑模型 (季节设定设为无)。

```
sales.ets(e=a, t=a, s=a, dgopt=m, dgraph=flts)
```

将使用 AAA (加法误差、加法趋势、加法季节) 模型对序列 SALES 进行平滑，并以 spool 对象显示输出，其中包含实际序列与分解序列 (即预测、趋势、水平与季节序列) 的多图。

```
sales.ets(e=a, t=a, s=a, level=level1, trend=trend1,  
          season=season1, dgopt=s, dgraph=flts)
```

将使用 AAA (加法误差、加法趋势、加法季节) 模型对序列 SALES 进行平滑，分别创建名为 level1、trend1 与 season1 的水平、趋势与季节分解序列，并在 spool 对象中以单一图形显示实际与分解图形。

```
tb3.ets(e=e, t=e, s=e, graph=c1)
```

将在不同误差、趋势与季节设定之中寻找最优模型，并在包含各可用模型预测与似然比较图的 spool 对象中呈现估计结果。

```
tb3.ets(e=a, t=e, s=e, amse, table=c1)
```

将使用平均均方误差计算搜索最优模型，并在包含预测与似然比较表的 spool 对象中显示估计结果。

交叉引用

参见《User's Guide I》第 89 页“ETS Exponential Smoothing”，其中讨论了指数平滑方法。

favar Var Methods

估计一个因子增广 VAR 设定。

语法

标准 FAVAR 的命令形式为

```
var_name.favar(options) lag_spec y_list @ exog_list @FAVARFACTORS factor_list
```

其中 lag_spec 为 VAR 滞后设定，y_list 为观测内生变量列表，exog_list 为外生变量列表，factor_list 为用于因子提取的序列列表。

BBE FAVAR 的命令形式为

```
var_name.favar(options) lag_spec y_list @ exog_list @FAVARSLLOW slow_list*  
@FAVAR-FAST fast_list*
```

其中 slow_list 与 fast_list 分别为慢速与快速因子变量设定。

若不需要外生回归变量，则以与标准 VAR 估计相同的方式输入空的外生字段。

选项

type=arg(default=bbe) FAVAR 类型：standard（使用@**FAVARFACTORS** 提供的因子变量估计标准 FAVAR）或 BBE（使用 @**FAVARSLLOW** 与 @**FAVARFAST** 估计 Bernanke-Boivin-Eliasz FAVAR，默认）。

demean 在提取因子前，对估计样本内的每个因子变量去均值。

sdize 在提取因子前，对估计样本内的每个因子变量标准化。

demeancross 在提取因子前，对因子变量横截面上的观测值去均值。

sdizecross 在提取因子前，对因子变量横截面上的观测值标准化。

mfmeth=arg 最大因子方法。支持的取值包括 schwert、ah、rootsize、

	size 与 user。
n=int	当 mfmethod=user 时，用户指定的最大因子数。
fsmethod=arg	因子选择方法。支持的取值包括 bn、ah、simple 与 user。
fsic_bn=arg	Bai-Ng 准则。支持的取值包括 ICP1、ICP2、ICP3、PCP1、PCP2、PCP3 与 AVG。
fsic_ah=arg	Ahn-Horenstein 准则。支持的取值包括 ER、GR 与 AVG。
fsic_simple=arg	简单因子选择准则。支持的取值包括 MIN、MAX 与 AVG。

mineigen=num	简单因子选择方法所用的最小特征值阈值。
cproport=num	简单因子选择方法所用的累积特征值比例阈值。
r=int	当 fsmethod=user 时，用户指定的保留因子数。

默认情况下，FAVAR 估计使用 type=bbe，不做去均值或标准化，mfmethod=user (n=1) 与 fsmethod=simple (fsic_simple=min)。简单方法的默认阈值为 mineigen=0 与 cproport=1。若选择 fsmethod=user，则使用 r=int 提供保留因子数，默认值为 r=1。默认的 Bai-Ng 准则为 fsic_bn=icp1，默认的 Ahn-Horenstein 准则为 fsic_ah=er。

示例

以下命令估计一个标准 FAVAR，含 4 阶滞后、2 个观测内生变量、一个常数项，以及从 XGROUP 中的变量提取的因子：

```
var1.favar(type=standard, fsmethod=user, r=3) 1 4 y1 y2 @ c
      @FAVARFACTORS xgroup
```

下一条命令估计一个 BBE FAVAR，含 4 阶滞后、一个常数项、慢速变量 XSLOW1 XSLOW2 XSLOW3 与快速变量 XFAST：

```
var2.favar(type=bbe, fsmethod=user, r=3) 1 4 policy @ c @FAVARSLow
      xslow1 xslow2 xslow3 @FAVARFAST xfast
```

以下命令估计一个标准 FAVAR，并使用 Bai-Ng 准则选择因子个数：

```
var3.favar(type=standard, mfmethod=schwert, fsmethod=bn,
      fsic_bn=avg, demeantime, sdizetime) 1 6 y1 y2 @ c @FAVARFACTORS
      x
```

交叉引用

参见《User's Guide II》第 355 页“Factor-Augmented VAR (FAVAR)”以获取更多讨论。

lagllama Series Procs

对底层序列执行 Lag-Llama 预测。

此命令需要随 EViews 安装的 EViews Python 库。

EViews 中的 lagllama 命令执行 Lag-Llama 自动预测。该命令是对 Python 库的封装，利用大语言模型类基础设施进行时间序列预测。

语法

`series_name.lagllama(options) forecast_name [quantile_prefix]`

选项

<code>nsamples=int</code>	设定执行的预测样本数（默认 200）。
<code>ctxtmult=int</code>	设定上下文乘数（默认 3）。
<code>quantiles="arg"</code>	指定预测分位数。默认预测 0.1 与 0.9 分位数。
<code>f=arg</code>	预测样本外填充行为：“actual”（以拟合变量的实际值填充预测样本外的观测值），或“na”（以缺失值填充预测样本外的观测值，默认）。
<code>seed=int</code>	为预测的随机抽取设定 Python 种子值。
<code>save=arg</code>	将所有预测抽取结果以名为 arg 的矩阵保存到工作文件中。

预测样本选项

预测样本从估计样本（即当前工作文件样本）之后的第一个观测值开始。预测终点由以下二者之一给定：

`forclen=int` 预测期数。
`forc="date"` 指定预测终点的日期。

若省略，则终点为工作文件范围的末尾。

示例

```
elecddm.lagllama(nsamples=100, ctxtmult=2, quantiles="0.2 0.8")
elecddm_f
```

使用 Lag-Llama 对 ELECDMD 序列生成预测，包含 100 次预测抽取、上下文乘数为 2，并保存 20% 与 80% 分位数。点预测将保存到 ELECDMD_F 序列中。

交叉引用

参见《User's Guide I》第 159 页“Lag-Llama”以获取更多讨论。

Lstm Series Procs

对底层序列执行长短期记忆（LSTM）估计与预测。

该过程构建滞后目标与可选特征输入，将学习样本按时间顺序划分为训练与验证部分，训练网络，并创建带可选诊断输出的预测序列。

语法

```
series_name.lstm(options) [forecast_name] @target target_name [@features
feature_list] [@hyperparams hp_options]
```

为输出预测序列输入一个名称。若省略 *forecast_name*，EViews 将基于目标序列名称与后缀 _LSTM_F 创建一个未使用的名称。

分组参数如下：

@ target target_name	标识 LSTM 设定的目标序列。正常情况下，target_name 与 series_name 为同一序列。
@ features feature_list	可选的特征序列或有效序列表达式的空格分隔列表。特征的滞后与标准化选项将应用于列出的每个特征。
@ hyperparams hp_options	可选的以逗号分隔的超参数赋值。固定值与调参控制在下方“超参数选项（Hyperparameter options）”中说明。

选项名称与关键字取值不区分大小写。对包含空格的选项值加引号，例如 dims_hidden="64 32"。

对于布尔型控制项，表中仅显示改变默认状态的写法：无前缀的名称启用默认关闭的功能，而 no 前缀则禁用默认开启的功能。

选项

数据与网络设定

normalize_series=arg	目标序列标准化。arg 可为 none、minmax（缩放至 [-1, 1]）、minmax01（缩放至 [0,1]）、stdize（标准化，默认）、robust、log 或 tanh。
lag=int	目标序列滞后设定（默认=0）。
nolagrange	仅使用选定的目标滞后；默认使用 0 到 lag 的全部范围。
nolagsasfeatures	将目标作为一个跨所包含滞后的特征来处理；默认将各滞后视为独立特征。
normalize_features=arg	特征标准化。arg 使用与 normalize_series 相同的关键字；默认值为 stdize。此选项仅在提供 @ features 时使用。
lag_features=int	应用于每个附加特征的滞后设定（默认=0）。
nolagrange_features	仅使用选定的特征滞后；默认使用 0 到 lag_features 的全部范围。
dims_hidden=arg	以空格分隔的正整数隐藏单元数列表，每个 LSTM 层对应一个（默认="1"）。例如，dims_hidden="32" 指定一层，而 dims_hidden="64 32" 指定两个堆叠层。
horizon_fcast=int	用于训练的目标步长（默认=1）。取值为 1 时训练单步向前模型；更大的值则针对所选预测水平训练模型。该过程为选定水平创建一个预测序列。
epochs=int	最大训练轮数（默认=500）。早停可能在达到此上限前结束训练。
ratio_train=num	按时间顺序分配给训练的学习样本比例（默认=0.80）。
ratio_validate=num	按时间顺序分配给验证的学习样本比例（默认=0.20）。训练与验证比例之和不得超过 1。若二者之和小于 1，则剩余的学习观测值未被使用。

标准化参数由学习样本的训练部分估计得到，随后原样应用于验证与预测观测值。更大的滞后设定会减少可用观测值的有效数量。

预测样本与样本内填充

学习样本为当前工作文件样本。预测样本从学习样本之后的第一个观测值开始。使用下列选项之一指定预测终点：

forcen=int	预测观测值个数。该值必须产生至少一个样本外预测观测值。
forc="date"	预测终点日期或观测值标签。若既未提供 forc 也未提供 forclen，则终点为工作文件范围的末尾。
f=arg	输出预测序列的样本内填充方式。arg 可为 na、actual、static（默认）或 dynamic。static 在可用时使用实际滞后目标值；dynamic 递归使用先前的预测值。

学习控制

loss=arg	损失函数。arg 可为 mse（默认）、mae、huber 或 ce。
weight_init=arg	权重初始化。arg 可为 zeros、ones、random、xavier（默认）、he、orthogonal 或 user。
regularization=arg	正则化类型。arg 可为 none（默认）、l1 或 l2。使用 @hyperparams 组中的 regularization_decay 设定罚项强度。
gradclip=arg	梯度裁剪方法。arg 可为 none、value 或 norm（默认）。
gradclip_norm=num	当 gradclip=norm 时使用的 L2 范数上限（默认=1.0）。
gradclip_lower=num	当 gradclip=value 时使用的逐元素裁剪下界（默认=-5）。
gradclip_upper=num	当 gradclip=value 时使用的逐元素裁剪上界（默认=5）。
randomize	使用随机化的连续学习子样本而非单一固定划分（默认=off）。
randomize_count=int	对非调参模型的随机化重复次数（默认=10）。仅在选择 randomize 时使用。
noearly_stop	禁用基于验证的早停（默认启用，并恢复表现最佳

	的权重)。
early_stop_warmup=int	在早停决策开始前所需的验证损失观测值个数 (默认=50)。
early_stop_patience=int	在停止前连续无合格验证损失改善的轮数 (默认=10)。
early_stop_mindelta_absolute=num	使耐心计数重置所需的最小绝对验证损失改善量 (默认=0.001)。
early_stop_mindelta_relative=num	使耐心计数重置所需的最小相对验证损失改善量 (默认=0.001)。满足任一正向阈值的变动均视为改善。

优化器与学习率调度器

gd=arg	梯度下降优化器。arg 可为 sgd、momentum、adagrad、rmsprop、adadelta、adam、adamw (默认)、nadam、radam 或 ranger。
rs=arg	学习率调度器。arg 可为 none、step、exponential、cosannealing、plateau (默认)、cycle 或 restarts。
rs_patience=int	在 plateau 调度器缩减前无验证改善的轮数 (默认=5)。

在 **@hyperparams** 组中设定学习率、优化器系数、调度器系数及其调参控制。

训练权重与输出

weights_page=name	用于导出训练权重或标识待加载权重的工作文件页面。导出时使用新的页面名，以避免与现有页面冲突。
noweights_export	禁止权重导出；当提供 weights_page 时默认启用导出。
weights_load	从 weights_page 加载已保存的权重及相关模型设定 (默认=off)。恢复的设定包括滞后、标准化、特征与层维度。
noplotloss	禁止输出训练与验证损失图 (默认生成)。
plotlrevo	生成学习率演化图。默认禁用。

plotlrevo	生成学习率演化图。默认禁用。
tableloss	生成按轮数的损失值表。默认禁用。
tabletuner	生成超参数调参器表。默认禁用，且仅在至少一个调参器处于激活状态时相关。

常规选项

prompt 在程序内部强制弹出 LSTM 对话框。

p 打印输出视图。

超参数选项

在 **@hyperparams** 之后提供超参数赋值；各赋值以逗号分隔。除下文注明的隐藏单元特例外，一般形式为：

@hyperparams root=value, root_tuner=key, root_min=value, root_max=value,
root_gridlen=int

root=value	设定固定的当前值。对于隐藏单元，当前层大小由 dims_hidden 设定；使用其余 dim_hidden_* 控制项来调整这些层大小。
root_tuner=key	调参器类型。key 可为 none（默认）、sequence（均匀间隔候选网格）或 random（随机候选网格）。
root_min=value	调参器为 sequence 或 random 时使用的搜索下界。
root_max=value	调参器为 sequence 或 random 时使用的搜索上界。
root_gridlen=int	为 root 生成的候选值个数（默认=10）。

支持的 root 及其默认值如下：

dim_hidden	隐藏单元。当前大小来自 dims_hidden；搜索最小 1、最大 100、网格长度 10。隐藏单元调参分别应用于每一层。
gd_rate_learn	优化器学习率。默认 0.001；搜索最小 0、最大 1、网格长度 10。gd_rate_learn_min 同时提供学习率调度器使用的绝对下界。
gd_momentum	动量或速度系数，由所选优化器使用。默认 0.9；搜索最小 0、最大 1、网格长度 10。
gd_decay1	第一优化器衰减或平滑系数。默认 0.9；搜索最小 0、最

dim_hidden	隐藏单元。当前大小来自 dims_hidden; 搜索最小 1、最大 100、网格长度 10。隐藏单元调参分别应用于每一层。大 1、网格长度 10。
gd_decay2	第二优化器衰减或平滑系数。默认 0.999; 搜索最小 0、最大 1、网格长度 10。
gd_decay3	额外的优化器特定衰减系数。默认 0.00025; 搜索最小 0、最大 1、网格长度 10。
rs_decay	step、exponential 与 plateau 调度器的缩减因子。默认 0.5; 搜索最小 0、最大 1、网格长度 10。
rs_cycle_len	余弦退火、循环调度与热重启的循环长度。默认 10; 搜索最小 10、最大 50、网格长度 10。
rs_cycle_step	step 与 exponential 调度器使用的衰减步长间隔。默认 10; 搜索最小 10、最大 50、网格长度 10。
rs_cycle_min	循环与基于余弦的调度器的学习率下界。默认 1×10^{-7} ; 搜索最小 0、最大 1、网格长度 10。
rs_cycle_max	循环与基于余弦的调度器的学习率上界。默认 0.1; 搜索最小 0、最大 1、网格长度 10。
regularization_decay	L1/L2 罚项强度。默认 0.1; 搜索最小 0、最大 1、网格长度 10。仅当 regularization=l2 时使用。

其他被调参的优化器、调度器与正则化参数对模型而言是全局的。当随机化与调参同时激活时，每个候选配置均使用为该候选生成的随机化设定进行评估。

示例

命令

```
elecddm.lstm(forclen=30, lag=14, dims_hidden="32") elecddm_lstm
@target elecddm
```

估计一个含 32 个隐藏单元的单层 LSTM，使用目标滞后 0 至 14，并在 ELECDMD_LSTM 中创建 30 个样本外预测。

命令

```
elecddm.lstm(forc="12/31/2015", lag=14, lag_features=0,
dims_hidden="64 32", plotlrevo, tableloss) elecddm_lstm @target
elecddm @features @month @quarter
```

使用两个堆叠的 LSTM 层，加入同期日历特征，设定预测终点，并请求学习率图与损失表。

命令

```
elecdmd.lstm(forclen=30, lag=14, dims_hidden="32", tabletuner)
  elecdmd_tuned @target elecdmd @features @month @quarter
  @hyperparams dim_hidden_tuner=sequence, dim_hidden_min=16,
  dim_hidden_max=64, dim_hidden_gridlen=4,
  gd_rate_learn_tuner=random, gd_rate_learn_min=0.0001,
  gd_rate_learn_max=0.01, gd_rate_learn_gridlen=8
```

为隐藏单元搜索一个四值顺序网格，为学习率搜索一个八值随机网格，并依据验证表现选择取值。

命令

```
elecdmd.lstm(forclen=30, lag=14, dims_hidden="32",
  weights_page=lstm_w) elecdmd_f @target elecdmd
elecdmd.lstm(forclen=30, weights_page=lstm_w, weights_load,
  noweight_export, weight_init=user) elecdmd_f_reuse @target
elecdmd
```

首先估计一个模型并将其训练权重导出到工作文件页面 LSTM_W，随后加载已保存的权重及相关模型设定，生成第二个预测序列。

交叉引用

参见《User's Guide I》第 137 页“Long Short-Term Memory (LSTM)”，其中讨论了模型构建、学习选项、诊断与实际设定指引。

knots Equation Views

查看由样条（spline）估计的方程的节点（knot）值表。

语法

eq_name.**knots**

示例

```
equation eq1.splinerreg elecdmd temp
show eq1.knots
```

估计一个方程，以序列 ELECDMD 为因变量，并使用 B 样条将 TEMP 序列建模为回归变量，随后显示所计算节点值的表。

交叉引用

参见《User's Guide II》第 261 页“Spline Regression”以获取更多讨论。

modavg Equation Methods

使用模型平均（model averaging）方法进行估计。

使用加权平均最小二乘（WALS）或贝叶斯模型平均（BMA）估计方程。

语法

```
eq_name.modavg(options) y x1 [x2 x3...] @ z1 z2...
```

指定因变量，后接在每个模型中都出现的变量列表，再后接“@”符号以及并非每个模型都必须出现的变量列表。

选项

method = arg 模型平均方法：“wals”表示加权平均最小二乘（默认），“bma”表示贝叶斯模型平均。

WALS 选项

walsprior=arg arg 可为“laplace”（默认）、“subbotin”或“weibull”。

BMA 选项

msp=arg	模型空间先验：“beta-binomial”（默认）或“binomial”。
a = number (default=1)	Beta(a,b) 中的第一个形状参数，其中 p 为任意可选回归变量的包含概率。仅当使用 beta-binomial 先验作为模型空间先验时适用。
b = number (default=1)	Beta(a,b) 中的第二个形状参数，其中 p 为任意可选回归变量的包含概率。仅当使用 beta-binomial 先验作为模型空间先验时适用。
p = number (default=0.5)	任意可选回归变量的先验包含概率。若各可选回归变量的先验包含概率不同，则将 p 设为开放单位区间中

的向量或值列表。

`gtype=arg`

用于设定 Zellner g 先验中 g 的方法类型。选项包括 “benchmark”（默认）、“uip”（单位信息先验，unit information prior）、“ric”（风险膨胀准则，risk inflation criterion）、“leb”（局部经验贝叶斯，local empirical Bayes）以及 “user”（用户指定值）。若选择 “user”，用户还必须设定 gvalue。

`gvalue=number`

由用户设定的 Zellner g 先验中的 g 值。使用此选项须将 gtype 设为 “user”。

示例

命令

```
equation eq1.modavg(method=bma,msp=binomial,p="0.25 0.25 0.5 0.5")
  y c @ x1 x2 x3 x4
```

执行贝叶斯模型平均，以 Y 为因变量，截距 C 作为唯一焦点回归变量，以及四个可选回归变量 X1、X2、X3 与 X4。模型空间先验设定为二项分布，X1、X2、X3 与 X4 的包含先验概率分别为 1/4、1/4、1/2 与 1/2。

交叉引用

参见《User's Guide II》第 289 页“Model Averaging”以获取更多讨论。

modelranking Equation Views

按后验概率对模型排序显示。

该视图仅适用于使用贝叶斯模型平均（BMA）估计的方程。

语法

`obj_name.modelranking(options)`

选项

`showtop=int`

要显示的排名靠前的模型数量。默认值为模型空间中的模型数量与 100 中的较小者。

示例

```
eq1.modelranking(showtop=20)
```

按后验概率对模型排序，并显示前 20 个模型的相关信息。

交叉引用

参见 *User's Guide II* 第 289 页的“Model Averaging”。

modelsize Equation Views

显示与模型规模相关的结果。

该视图仅适用于使用贝叶斯模型平均（BMA）估计的方程。

语法

obj_name.**modelsize**

交叉引用

参见 *User's Guide II* 第 289 页的“Model Averaging”。

nprophet Series Procs

对底层序列执行 NeuralProphet 预测。

Prophet 需要随 EViews 安装的 NeuralProphet Python 库。

语法

series_name.**nprophet**(options) forecast_name [quantile_prefix]

输入预测序列的名称，以及（可选）分位数预测的名称前缀。

选项

growth=arg 设置增长类型。arg 可为 “linear”（默认）或 “log”。

season=arg 设置季节性类型。arg 可为 “add”（加法，默认）或 “mult”（乘法）。

trendreg=arg 设置趋势正则化参数。

growth=arg	设置增长类型。arg 可为 "linear"（默认）或 "log"。
seasreg=arg	设置季节性正则化参数。
quantiles="arg"	指定预测分位数。默认预测 0.1 与 0.9 分位数。
arlags=int	设置在 AR Net 分量中估计的自回归滞后阶数。
arreg=arg	设置 AR 分量正则化参数。
arlayers="list"	指定 AR Net 分量中隐藏层的层数与大小。"list" 应为以逗号分隔的层大小列表。例如，要包含两个大小分别为 64 和 32 的层，输入 arlayers= "64, 32"。
epochs=int	设置 AR-Net 分量中的迭代轮数（epochs）。
learnrate=arg	设置 AR-Net 分量中的学习率。
f=arg	样本外预测填充行为："actual"（用拟合变量的实际值填充预测样本之外的观测值）、"na"（用缺失值填充预测样本之外的观测值，默认），或 "forc"（用样本内预测值填充训练数据观测值，其余数据用 NA 填充）。
userregs="list"	包含用户自定义外生变量。"list" 应为以空格分隔的序列或 EViews 序列表达式列表，并用引号括起。
userregmodes="list"	设置用户自定义外生变量的模式。"list" 应为以空格分隔的 "mult" 或 "add" 列表，其中列表项数量需与 userregs 选项中的项数量一致。
userregpriors="list"	设置用户自定义外生变量的先验尺度。"list" 应为以空格分隔的尺度列表，其中列表项数量需与 userregs 选项中的项数量一致。
countryhol=arg	包含 Prophet 内置的某国假日集之一。arg 应为该国的国家代码。
countryreg=arg	设置该国假日的正则化参数。
countrylow=int	设置该国假日的下窗口。要指定假日之前的天数，该整数应为负数。
countryhigh=int	设置该国假日的上窗口。
countrymode=arg	设置该国假日的模式。arg 应为 "mult" 或 "add"（默认）。
userhols="list"	包含用户自定义事件或假日。"list" 应为以空格分隔的日期集或 EViews 假日关键字。

`userhollows="list"` 指定用户自定义事件或假日的下窗口集。“list”应为以空格分隔的整数集，用于指定将用户假日偏移多少天以设定假日期间的起始。注意：要指定假日之前的天数，该整数应为负数。

`userholhighs="list"` 指定用户自定义事件或假日的上窗口集。“list”应为以空格分隔的整数集，用于指定将用户假日偏移多少天以设定用户假日的结束。

`userholmodes="list"` 设置用户假日模式。“list”应为以空格分隔的“mult”或“add”列表，其中列表项数量需与 `userhols` 选项中的项数量一致。

`userholpriors="list"` 设置用户假日的先验尺度。“list”应为以空格分隔的尺度列表，其中列表项数量需与 `userhols` 选项中的项数量一致。

预测样本选项

预测样本将从紧随估计样本（当前工作文件样本）之后的第一个观测值开始。预测终点由以下二者之一给出：

`forcen=int` 预测的期数。

`forc="date"` 指定预测终点的日期。

若省略，终点将为工作文件范围的末尾。

示例

```
elec dmd.nprophet(f=actual, countryhol=UK, userregs="gdp
    @log(weather") elec_f elec_f_q
```

这将在 ELEC DMD 序列上执行 NeuralProphet 例程，对预测样本之外的观测值插入实际值，并包含 NeuralProphet 为英国（United Kingdom）内置的假日集，同时以 GDP 与 WEATHER 的对数作为外生回归变量。预测值将存储在新序列 ELEC_F 中，10% 与 90% 分位数分别存储于 ELEC_F_Q_1 与 ELEC_F_Q_9。

交叉引用

参见 *User's Guide I* 第 116 页的“NeuralProphet”以了解更多讨论。

`partialeffects` Equation Views

查看使用样条估计的方程的偏效应图形。

语法

eq_name.**partialeffects**

示例

```
equation eq1.splinerreg elecdmd temp  
show eq1.partialeffects
```

以序列 ELECDMD 为因变量估计方程，并使用 B 样条将 TEMP 序列建模为回归变量，随后显示偏效应图形。

交叉引用

参见 *User's Guide II* 第 261 页的“Spline Regression”以了解更多讨论。

placebo Equation Views

查看时间回归断点模型的安慰剂（placebo）图。

语法

eq_name.placebo [dates]

若未提供日期，EViews 将自动确定一组合适的安慰剂日期。

示例

```
equation eq1.rdit log(conceptions) @ 2007m7  
eq1.placebo 2006m3 2006m6 2008m3 2008m6
```

估计一个时间回归断点模型，以 CONCEPTIONS 的对数为结果变量、2007 年 7 月为干预日期，随后显示安慰剂图，安慰剂日期为 2006 年 3 月、2006 年 6 月、2008 年 3 月与 2008 年 6 月。

交叉引用

参见 *User's Guide II* 第 273 页的“Regression Discontinuity in Time (RDiT)”以了解更多讨论。

Var Procs

将 BVAR 或 BTVCVAR 的后验抽样 (posterior draws) 放入当前工作文件。

BVAR 说明

对于 BVAR，后验抽样被复制、向量化，然后作为一个或多个新建矩阵对象中的行进行存储。命名矩阵对象会被添加到工作文件中。

每个矩阵对象代表一个不同的参数块。参数块包括：系数矩阵 B、误差协方差矩阵 S，以及未知先验超参数 g。先验类型决定某一参数块是否相关：

- B 在所有先验类型下均相关。
- S 在所有先验类型下均相关，Litterman/Minnesota 先验除外。
- g 仅在 GLP 下相关。

如果后验抽样不可用，系统可能会提示您重新估计 BVAR。

版本兼容性说明：该过程在 EViews 15 中引入，替代了 EViews 11 至 14 中的 `var::drawcoefs` 与 `var::drawrescov`。

BTVCVAR 说明

对于 BTVCVAR，后验抽样被复制并作为序列对象存储在一个新页面中。每个序列保存某一特定参数的抽样。

如果抽样不可用，EViews 可能会提示您运行后验采样器。

语法

BVAR 语法

```
var_name.postdraws(options) [matrix_name_coef] [matrix-  
_name_ecov] [matrix_name_hype]
```

其中 *matrix_name_coef*、* *matrix_name_ecov* 与 *matrix_name_hype** 用于命名包含 B、S 与 g 抽样的矩阵对象。

BTVCVAR 语法

```
var_name.postdraws new_page_name
```

选项

BVAR 选项

<code>bylag</code>	若包含此选项，滞后系数将按滞后而非按变量分组。
<code>dropunstable</code>	若包含此选项，来自参数空间不稳定区域的抽样将被丢弃。
<code>burn=number</code>	预烧比例。注意 <code>number</code> 为开单位区间 (0, 1) 中的一个值。该选项仅在与马尔可夫链蒙特卡洛 (MCMC) 方法获得的抽样相关时适用。默认值为估计所用的预烧值。

示例

若 MYBVAR 是一个使用 Litterman/Minnesota 先验之外其他先验的贝叶斯 VAR，则

```
mybvar.postdraws(dropunstable) coefs ecovs
```

创建两个命名矩阵 COEFS 与 ECOVS。系数抽样作为行存储在 COEFS 中，误差协方差矩阵抽样作为行存储在 ECOVS 中。来自参数空间不稳定区域的抽样被排除在 COEFS 与 ECOVS 之外。

交叉引用

参见 *User's Guide II* 第 369 页的“Bayesian VAR Models”以了解更多讨论。

`postmarginals` Equation Views

显示参数后验分布的单变量边际 (marginals)。

该视图仅适用于使用贝叶斯模型平均 (BMA) 估计的方程。

该视图是一个包含三个部分的 `spool`。第一部分显示必选回归变量上系数的后验密度与后验均值。第二部分显示可选回归变量上系数的后验密度、后验均值与后验排除概率。第三部分显示误差标准差与误差方差的后验密度与后验均值。

语法

`eq_name.postmarginals`

交叉引用

参见 *User's Guide II* 第 289 页的“Model Averaging”以了解更多讨论。

prophet Series Procs

对底层序列执行 Facebook 的 Prophet 预测。

Prophet 需要随 EViews 安装的 Prophet Python 库。

语法

`series_name.prophet(options) forecast_name [lower_name upper_name]`

输入预测序列的名称，以及（可选）预测下界与上界的名称。

选项

<code>cpscale=number</code>	设置变点先验尺度。
<code>growth=arg</code>	设置增长类型。 <code>arg</code> 可为 “linear”（默认）或 “log”。
<code>season=arg</code>	设置季节性类型。 <code>arg</code> 可为 “add”（加法，默认）或 “mult”（乘法）。
<code>f=arg</code>	样本外预测填充行为：“actual”（用拟合变量的实际值填充预测样本之外的观测值）、“na”（用缺失值填充预测样本之外的观测值，默认），或 “forc”（用样本内预测值填充训练数据观测值，其余数据用 NA 填充）。
<code>bounds=arg</code>	设置下界与上界的区间宽度。
<code>userregs="list"</code>	包含用户自定义外生变量。“list” 应为以空格分隔的序列或 EViews 序列表达式列表，并用引号括起。
<code>userregmodes="list"</code>	设置用户自定义外生变量的模式。“list” 应为以空格分隔的 “mult” 或 “add” 列表，其中列表项数量需与 <code>userregs</code> 选项中的项数量一致。
<code>userregpriors="list"</code>	设置用户自定义外生变量的先验尺度。“list” 应为以空格分隔的尺度列表，其中列表项数量需与 <code>userregs</code> 选项中的项数量一致。
<code>countryhol=arg</code>	包含 Prophet 内置的某国假日集之一。 <code>arg</code> 应为该国的国家代码。

countryhol=arg	包含 Prophet 内置的某国假日集之一。arg 应为该国的国家代码。
countryprior=arg	设置该国假日的先验尺度。
userhols="list"	包含用户自定义事件或假日。“list”应为以空格分隔的日期集或 EViews 假日关键字。
userhollows="list"	指定用户自定义事件或假日的下窗口集。“list”应为以空格分隔的整数集，用于指定将用户假日偏移多少天以设定假日期间的起始。注意：要指定假日之前的天数，该整数应为负数。
userholhighs="list"	指定用户自定义事件或假日的上窗口集。“list”应为以空格分隔的整数集，用于指定将用户假日偏移多少天以设定用户假日的结束。
userholmodes="list"	设置用户假日模式。“list”应为以空格分隔的“mult”或“add”列表，其中列表项数量需与 userhols 选项中的项数量一致。
userholpriors="list"	设置用户假日的先验尺度。“list”应为以空格分隔的尺度列表，其中列表项数量需与 userhols 选项中的项数量一致。
prompt	强制在程序中弹出对话框。

预测样本选项

预测样本将从紧随估计样本（当前工作文件样本）之后的第一个观测值开始。预测终点由以下二者之一给出：

forclen=int 预测的期数。

forc="date" 指定预测终点的日期。

若省略，终点将为工作文件范围的末尾。

示例

```
elecdmd.prophet(f=actual) elec_f elec_l elec_u
```

这将在 ELECDMD 序列上执行 Facebook 的 Prophet 例程，对预测样本之外的观测值插入实际值。预测值将存储在新序列 ELEC_F 中，预测的下界与上界分别存储于 ELEC_L 与 ELEC_U。

交叉引用

参见 *User's Guide I* 第 111 页的“Facebook Prophet Forecasting”以了解更多讨论。

randomforest Series Procs

执行随机森林风格的回归预测

randomforest 需要随 EViews 安装的 EViews Python 库。

EViews 中的 randomforest 命令允许用户执行随机森林回归预测。该命令是对 SciKit Learn 库中随机森林实现的封装 (wrapper) 。

语法

`series_name.randomforest(options) forecast_name [exog_variables]`

选项

<code>model=arg</code>	指定所使用的随机森林模型类型。arg 可为 “rf” (随机森林, 默认)、 “dt” (决策树)、 “et” (极端随机树), 或 “gb” (梯度提升) 。
<code>lags=int</code>	设置模型中包含的滞后阶数。
<code>nodropseries</code>	不从外生变量 / 因子列表中丢弃含有内部缺失值的序列。
<code>f=arg</code>	样本外预测填充行为: “actual” (用拟合变量的实际值填充预测样本之外的观测值)、 “na” (用缺失值填充预测样本之外的观测值, 默认), 或 “forc” (用样本内预测值填充训练数据观测值, 其余数据用 NA 填充) 。
<code>trees=int</code>	设置森林中使用的树的数量 (若 <code>model=dt</code> 则不适用) 。
<code>depth=int</code>	设置树的最大深度。
<code>leafsize=int</code>	设置叶节点的最小大小。
<code>gain=num</code>	设置发生分裂所需的最小增益。
<code>nobootstrap</code>	禁用自助抽样 (bootstrap) (若 <code>model=dt</code> 则不适用) 。
<code>seed=int</code>	设置随机数生成的种子。
<code>learnrate=num</code>	设置梯度提升模型的学习率。
<code>subsample=num</code>	设置梯度提升模型的子样本比例。

subsample=num 设置梯度提升模型的子样本比例。
sharp=arg 指定包含输出 Sharp 值的工作文件矩阵对象的名称
import=arg 指定包含输出 Importance 值的工作文件矩阵对象的名称

预测样本选项

预测样本将从紧随估计样本（当前工作文件样本）之后的第一个观测值开始。预测终点由以下二者之一给出：

forcen=int 预测的期数。
forc="date" 指定预测终点的日期。

若省略，终点将为工作文件范围的末尾。

示例

```
elec dmd.randomforest(trees=100, depth=10, lags=2) elec dmd_f  
                      features_group
```

基于组 FEATURES_GROUP 中的外生变量，使用 100 棵树、最大深度 10，并在模型中包含 2 阶滞后，对序列 ELECDMD 执行随机森林预测。结果将存储在 ELECDMD_F 序列中。

交叉引用

参见 *User's Guide I* 第 163 页的“Random Forest”以了解更多讨论。

rditEquation Methods

估计一个时间回归断点模型。

语法

```
eq_name.rdit(options) y [z1 z2...] @ idate
```

先列出结果变量，随后是任意可选协变量，接着是一个 @ 符号加上干预/断点日期。

选项

polyest=int 指定估计多项式阶数 p。int 可为 0、1（默认）或 2。
polybias=int 指定偏差校正多项式阶数 q。int 可为 1、2（默认）或 3，且必须大于估计多项式阶数。

polyest=int	指定估计多项式阶数 p。int 可为 0、1（默认）或 2。
kern=arg	指定估计过程中使用的核（kernel）形状。arg 可为“Triangular”（默认）、“Epanechnikov”或“Uniform”。
bwselect=arg	指定带宽选择方法。arg 可为“mserd”（MSE 最优（RD），默认）、“msetwo”（MSE 最优（双侧））、“msesum”（MSE 最优（求和））、“msecomb1”（MSE 最优（组合 1））、“msecomb2”（MSE 最优（组合 2））、“cerrd”（CER 最优（RD））、“certwo”（CER 最优（双侧））、“cersum”（CER 最优（求和））、“cercomb1”（CER 最优（组合 1））、“cercomb2”（CER 最优（组合 2））、“user”（用户指定）。
userbw=arg	若 bwselect 选项设为“user”，此选项指定用户带宽。arg 应为用引号括起的、以逗号分隔的 4 个值列表（hl, hr, bl, br）。
donutl=int	设置左侧甜甜圈大小（默认为 0）。
donutr=int	设置右侧甜甜圈大小（默认为 0）。
cov=arg	指定所使用的协方差方法。cov 可为“ordinary”、“white”、“hc”（默认）或“hac”。
hctype=arg	指定 HC 协方差方法。仅在使用 cov=hc 选项时适用。arg 可取“hc1”、“hc2”（默认）或“hc3”。
covlag=arg(default=1)	白化（whitening）滞后设定：整数（用户指定滞后值），或“a”（自动选择）。
covinfosel=arg(default="aic")	自动选择的信息准则：“aic”（Akaike，赤池）、“sic”（Schwarz，施瓦茨）、“hqc”（Hannan-Quinn，汉南-奎因）（若“lag=a”）。
covmaxlag=integer	自动选择的最大滞后长度（可选）（若“lag=a”）。默认为基于观测值的最大值 $T^{(1/3)}$ 。
covkern=arg(default="bart")	核（kernel）形状：“none”（无核）、“bart”（Bartlett，巴特利特，默认）、“bohman”（Bohman）、“daniell”（Daniel）、“parzen”（Parzen）、“parzriesz”（Parzen-Riesz）、“parzgeo”（Parzen-Geometric）、“parzcauchy”（Parzen-Cauchy）、“quadspec”（Quadratic Spectral，二次谱）、“trunc”（Truncated，截断）、“thamm”（Tukey-Hamming）、“thann”（Tukey-Hanning）、

`polyest=int` 指定估计多项式阶数 `p`。`int` 可为 0、1（默认）或 2。
“`tparz`”（Tukey-Parzen）。

`covbw=arg(default="fixednw")` 核带宽：“`fixednw`”（Newey-West 固定）、“`andrews`”（Andrews 自动）、“`neweywest`”（Newey-West 自动）、数字（用户指定带宽）。

`covnwlag=integer` 用于非参数核带宽选择的 Newey-West 滞后选择参数（若 “`covbw=neweywest`”）。

`covbwint` 使用带宽的整数部分。

示例

```
equation eq1.rdit log(conceptions) @ 2007m07
```

估计一个时间回归断点模型，以 CONCEPTIONS 的对数为结果变量，2007 年 7 月为干预日期。

```
equation eq2.rdit(polybias=3, bwselect=user, userbw="6.5, 6.5,  
8.0, 8.5") log(conceptions) @ 2007m7
```

估计同一模型，但将用于偏差校正的多项式阶数提高到 3，并将带宽固定为 `hl=6.5`、`hr=6.5`、`bl=8.0`、`br=8.5`。

交叉引用

参见 *User's Guide II* 第 273 页的“Regression Discontinuity in Time (RDIT)”以了解更多讨论。

`splinedecom` Series Procs

使用样条分解对序列进行季节调整。

语法

```
series_name.command(options) [seasonal_spec] seasadj_name [factor_name  
trend_name]
```

应在 `splinedecom` 关键字后跟随一个可选的季节性设定，再接季节调整后序列的名称。可选地，您还可以提供输出因子序列与趋势序列的名称。

选项

<code>tform=arg</code>	在执行样条分解之前对底层数据进行变换。 <code>arg</code> 可为“none”（默认）、“box-cox”、“yeo-johnson”、“log”、“square-root”（平方根）、“arcsin”（反正弦）、“auto”。
<code>tformlambda=num</code>	设置 Box-Cox 与 Yeo-Johnson 变换的幂水平。
<code>trmethod=arg</code>	指定分解中趋势分量所用的样条方法。 <code>arg</code> 可为“bspline”、“polyspline”、“pspline”（默认）。
<code>trdegree=int</code>	设置趋势样条分量的阶数。默认值为 3。
<code>trpenalty=num</code>	设置趋势分量的 p 样条罚项（仅当 <code>trmethod</code> 设为“pspline”时适用）。默认值为 0.8。
<code>noseas</code>	不包含季节性分量。默认情况下 EViews 会确定一个合适的季节性分量纳入分解。
<code>noseastest</code>	不执行季节性检验。
<code>noseasplot</code>	不显示季节性图形。

季节设定

使用 **@seas** 关键字并后接选项，以指定一个季节性周期样条分量。通过多次使用 **@seas** 设定，可以包含多个季节性分量。

季节选项

`degree=int` 设置周期样条的阶数。
`cycle=int` 设置周期样条的周期性。

示例

```
houstnsa.splinedecomp houst houst_f houst_tr
```

对序列 HOUSTNSA 执行样条分解季节调整，生成名为 HOUST 的季节调整后序列，季节因子存入序列 HOUST_F，趋势分量存入 HOUST_TR。

```
houstnsa.splinedecomp(tform=yeo-johnson, tformlambda=3) houst  
houst_f houst_tr
```

执行相同的分解，但使用 $\lambda=3$ 的 Yeo-Johnson 变换对 HOUSTNSA 序列进行变换。

```
sp500.splinedecomp(trdegree=2) @seas(degree=2, cycle=366)
    @seas(degree=3, cycle=7) spsa
```

对序列 SP500 进行季节调整，将趋势分量的样条阶数设为 2，并包含两个季节性分量：一个阶数为 2、周期为 366，另一个阶数为 3、周期为 7。

交叉引用

参见 *User's Guide I* 第 250 页的“Spline Decomposition Seasonal Adjustment”以了解更多讨论。

splnereg Equation Methods

使用非参数样条估计方程。

语法

```
eq_name.splnereg(options) y x1 [x2 x3...] [@ z1 z2...]
```

列出因变量，随后是用样条建模的非参数外生回归变量列表。可选地，您可以使用“@”符号后接线性外生回归变量列表。设定中不应包含常数项。

选项

method=arg	选择样条类型。arg 可为“bspline”（默认）、“polyspline”或“pspline”。
degree=int	指定样条的阶数。
knotsloc=arg	指定节点（knot）位置的选择方法。arg 可为“auto”（在“quantile”与“uniform”之间自动选择）、“quantile”（底层变量的分位数，默认）、“uniform”（底层变量最小值与最大值之间的均匀分布），或以空格分隔的值列表，或是包含节点值的工作文件向量对象名称。
knotsnum=arg	指定内部节点的数量。arg 可为“auto”（自动选择，默认），或固定整数。
maxknots=arg	自动选择时内部节点的最大数量。
boundaryadjust=num	指定将边界节点扩展到底层变量最小/最大值之外的比例因子。默认值为 0。