

02427 lecture 8: Hierarchical forecast reconciliation

Jan Kloppenborg Møller

DTU-Compute - Technical University of Denmark
jkmo@dtu.dk

Advanced time series analysis – 2025

Outline

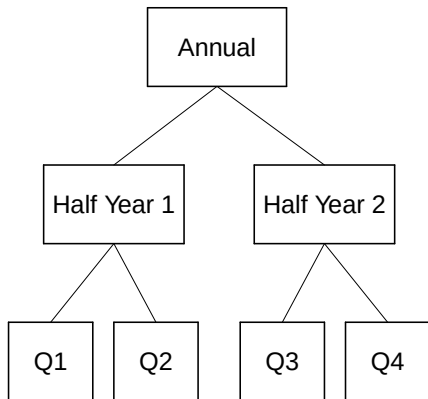
- 1 Introduction
- 2 Heat load forecasting
- 3 Likelihood inference
- 4 The general linear model
- 5 Conclusion

Outline

- 1 Introduction
- 2 Heat load forecasting
- 3 Likelihood inference
- 4 The general linear model
- 5 Conclusion

Motivation

- Models on different aggregation levels
- Models on each level may not agree
- Reconciliation ensure consistent forecasts
- Reconciliation often improve forecast accuracy on all levels



Mathematical formulation

The reconciled forecast, $\tilde{\mathbf{y}}$, is found as the solution to the optimization problem

$$\text{minimize} \quad (\hat{\mathbf{y}} - \tilde{\mathbf{y}})^T \boldsymbol{\Sigma}^{-1} (\hat{\mathbf{y}} - \tilde{\mathbf{y}})$$

subject to the coherence constraints $\tilde{\mathbf{y}} = \mathbf{S}\mathbf{P}\tilde{\mathbf{y}}$.

In a Gaussian setting the model can be formulated as

$$\hat{\mathbf{y}} = \mathbf{S}\tilde{\mathbf{y}} + \boldsymbol{\epsilon}; \quad \boldsymbol{\epsilon} \sim N(\mathbf{0}, \boldsymbol{\Sigma})$$

with an equivalent solution. Note that the dimension of $\tilde{\mathbf{y}}$ is different in the two settings.

Reconciliation

The reconciled forecast is calculated by

$$\tilde{\mathbf{y}} = (\mathbf{S}^T \mathbf{\Sigma}^{-1} \mathbf{S})^{-1} \mathbf{S}^T \mathbf{\Sigma}^{-1} \hat{\mathbf{y}}$$

where:

- $\tilde{\mathbf{y}}$: reconciled forecast
- $\mathbf{S}$: the summation matrix
- $\mathbf{\Sigma}$: a variance-covariance matrix
- $\hat{\mathbf{y}}$: the base forecast

$$\mathbf{S} = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}$$

Regression setting

$$\hat{\mathbf{y}} = \mathbf{S} \tilde{\mathbf{y}} + \boldsymbol{\epsilon}; \quad \boldsymbol{\epsilon} \sim N(\mathbf{0}, \mathbf{\Sigma})$$

Choosing the correct variance

The “only” unknown part is the variance-covariance matrix Σ , the optimal choice is (Wickramasuriya et al. 2019)

$$\Sigma = V[\mathbf{y} - \hat{\mathbf{y}}].$$

Hence a natural choice is to use the observed counturpart, i.e.

$$\hat{\Sigma} = \frac{1}{T} \sum_{t=1}^T (\mathbf{y}_t - \hat{\mathbf{y}}_t)(\mathbf{y}_t - \hat{\mathbf{y}}_t)^T$$

In realistic settings this matrix is usually high dimensional and ill-conditioned, and hence other versions are usually used.

In sample properties

With $\hat{\Sigma}$ as above, then:

For any column (and row) selection, $I \subseteq \{1, 2, \dots, n\}$, of $\hat{\mathbf{Y}}$ (and $\mathbf{S}$) then

$$\mathbf{S}_I (\mathbf{Y} - \tilde{\mathbf{Y}})^\top (\mathbf{Y} - \tilde{\mathbf{Y}}) \mathbf{S}_I^\top \leq (\mathbf{Y} \mathbf{S}_I^\top - \hat{\mathbf{Y}}_I)^\top (\mathbf{Y} \mathbf{S}_I^\top - \hat{\mathbf{Y}}_I), \quad (1)$$

and hence for the in-sample variance of the residuals are reduced on any level.

Choosing the correct variance

Different options for Σ :

- $\Sigma = I$
- Variance scaling (proportional to volume of level)
- Use observed variance-covariance of base forecast error
- Ignore cross level correlation
- Use shrinkage on the observed variance-covariance (usually preferred).
I.e.

$$\hat{\Sigma}_s = \lambda \hat{\Sigma} + (1 - \lambda) \text{diag} \hat{\Sigma}$$

Further the optimal λ is $\hat{\lambda} = \frac{\sum_{i \neq j} \hat{v} \hat{r}_{ij}}{r_{ij}^2}$, with r_{ij} elements of the correlation matrix.

Choosing the correct variance – simple examples

Some very simple examples are

$$\Sigma_{struc} = \text{diag}(4, 2, 2, 1, 1, 1, 1)$$

$$\Sigma_{svar} = \text{diag}(\sigma_A^2, \sigma_H^2, \sigma_H^2, \sigma_Q^2, \sigma_Q^2, \sigma_Q^2, \sigma_Q^2)$$

$$\Sigma_{hvar} = \text{diag}(\sigma_A^2, \sigma_{H_1}^2, \sigma_{H_2}^2, \sigma_{Q_1}^2, \sigma_{Q_2}^2, \sigma_{Q_3}^2, \sigma_{Q_4}^2)$$

These all ignore cross covariance, and as we will see these might be important to obtain optimal reconciliation.

Choosing the correct variance – block structures

Also block diagonal structures have been suggested

$$\Sigma_{acov} = \begin{bmatrix} \sigma_A^2 & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \Sigma_H & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \Sigma_Q \end{bmatrix}$$

which might further include some parametrized structures in the block diagonals. Here cross level covariance are ignores.

In conclusion the shrinkage estimator is usually preferred (Nystrup et al 2019).

Estimation of the precision matrix – dimensionally reduction

First observe that, while, it is usually not reasonable to assume that the variance–covariance matrix is sparse. It is often reasonable to assume that the precision matrix (Σ^{-1}) is (e.g. AR(p) models).

Using this observation formulate the precision as

$$\Sigma^{-1} = \Lambda^{-1/2} \Theta \Lambda^{-1/2}$$

where Λ is a diagonal matrix with diagonal elements equal the marginal variances.

Estimation of the precision matrix – dimensionally reduction

Dimensionality reduction can be obtained by e.g.

- GLASSO, i.e. maximize the penalized log-likelihood:

$$\log |\Theta| - \text{tr}(\mathbf{R}\Theta) - \lambda \sum_{i \neq j} |\Theta_{ij}|$$
- λ is a regularization parameter and larger λ lead to simpler inverse correlation matrices.

Another option is spectral scaling (Nystrup et al 2021), where Θ is chosen as

$$\Theta = (\mathbf{W}\mathbf{A}\mathbf{W}^T + \sigma^2\mathbf{I})^{-1}$$

where $\mathbf{W}$ is the first n_e eigen vectors of the correlation matrix, and $\mathbf{A}$ is related to the eigenvalues. $\sigma^2\mathbf{I}$ is a regularization/shrinkage.

Example

Example: Assume that half year levels generated by

$$Y_t = \phi_1 Y_{t-1} + \phi_2 Y_{t-2} + \epsilon_t,$$

AR(1) models half-year levels and annual levels, i.e. the models

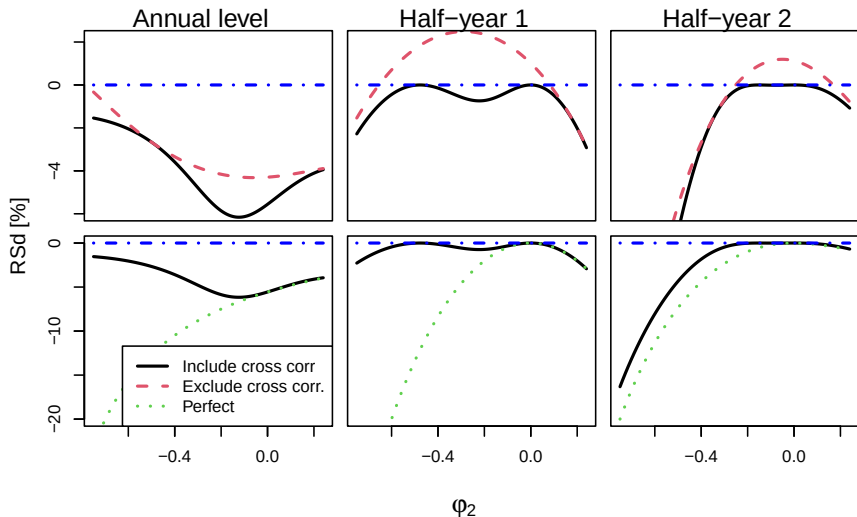
$$\begin{aligned} y_t^A &= \phi_1^A y_{t-1}^A + \epsilon_t^A \\ y_t^H &= \phi_1^H y_{t-1}^H + \epsilon_t^H. \end{aligned}$$

Full setup

$$\mathbf{y}_{2t+2|2t} = \begin{bmatrix} y_{2t+2|2t}^A \\ y_{2t+1|2t}^H \\ y_{2t+2|2t}^H \end{bmatrix}; \quad \hat{\mathbf{y}}_{2t+2|2t} = \begin{bmatrix} \hat{y}_{2t+2|2t}^A \\ \hat{y}_{2t+1|2t}^H \\ \hat{y}_{2t+2|2t}^H \end{bmatrix}; \quad \mathbf{S} = \begin{bmatrix} 1 & 1 \\ 1 & 0 \\ 0 & 1 \end{bmatrix}$$

and e.g. $\Sigma = \text{Var}[\mathbf{y}_{2t+2} - \hat{\mathbf{y}}_{2t+2}]$ can be calculated explicitly.

Choosing the correct variance-covariance



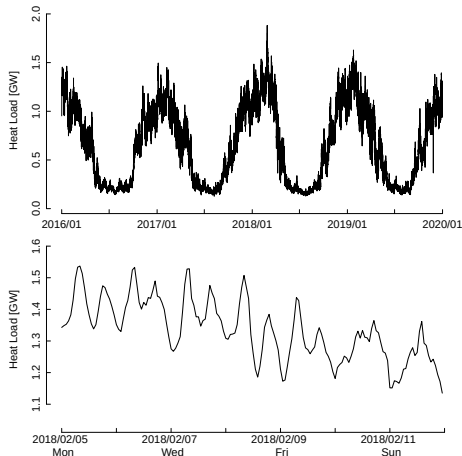
$$\phi_1 = 0.75$$

Outline

- 1 Introduction
- 2 Heat load forecasting**
- 3 Likelihood inference
- 4 The general linear model
- 5 Conclusion

Data

- District heating from an area of greater Copenhagen
- Clear annual and diurnal variation
- State of the art commercial hourly forecast (1-24 hours ahead)
- 2016 used for initialization



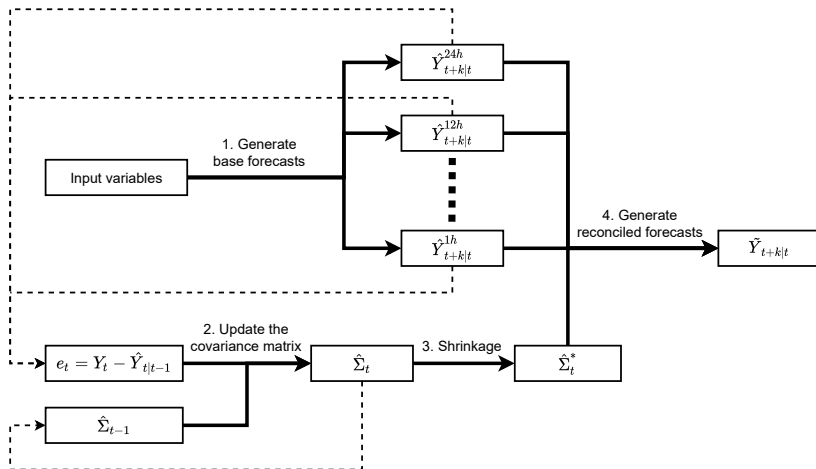
From Bergsteinsson et al. (2021).

A case study

- 1 hour level used for evaluation
- All other levels modeled in the study
- 2, 3, 4, 6, 8, 12, and 24 hours forecast:
 - Recursive Least Square
 - Forecast of ambient temperature
 - Diurnal variation
 - Auto-regressive parts
- Σ_t (60 by 60) estimated using the full variance covariance with shrinkage, and recursive updating.

Work flow of modeling

5. Compute forecast error when observations become available



From Bergsteinsson et al. (2021).

Some results (RRMSE)

	2017–2019			
	Base RMSE	Expanding Window	Rolling Window	Exponential Smoothing
Daily	0.5960	-23.75	-22.49	-23.93
Twelve-hourly	0.3516	-24.08	-22.83	-24.2
Eight-hourly	0.3538	-43.51	-42.72	-43.69
Six-hourly	0.2876	-44.64	-43.75	-44.76
Four-hourly	0.1765	-36.06	-35.19	-36.37
Three-hourly	0.1334	-33.03	-32.05	-33.26
Two-hourly	0.0884	-30.09	-29.07	-30.36
Hourly	0.0383	-14.75	-13.46	-15.07

Adapted from Bergsteinsson et al. (2021).

Outline

- 1 Introduction
- 2 Heat load forecasting
- 3 Likelihood inference**
- 4 The general linear model
- 5 Conclusion

Motivation / Aim

The starting point

$$\tilde{\mathbf{y}} = (\mathbf{S}^T \mathbf{\Sigma}^{-1} \mathbf{S})^{-1} \mathbf{S}^T \mathbf{\Sigma}^{-1} \hat{\mathbf{y}}$$

Comments and aim

- Observation appear only through the variance-covariance matrix $\mathbf{\Sigma}$
- The estimation of $\mathbf{\Sigma}$ include a large number of parameters
- Formulate a parameterized model for obtaining $\mathbf{\Sigma}$
- Reduce dimension of the parameter space by well-known likelihood techniques

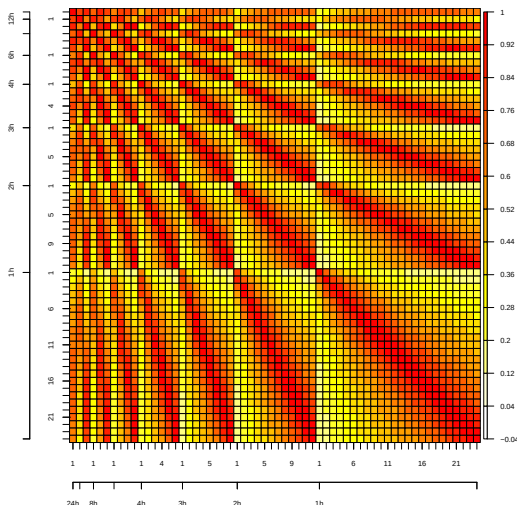
Modeling of variance-covariance matrix

Σ (some challenges):

- High-dimensional
- High correlations

Some suggestions

- Parametric models for the correlation
- Use statistical methods for reduction



Graphics: Møller et al. (2023)

A parametric model

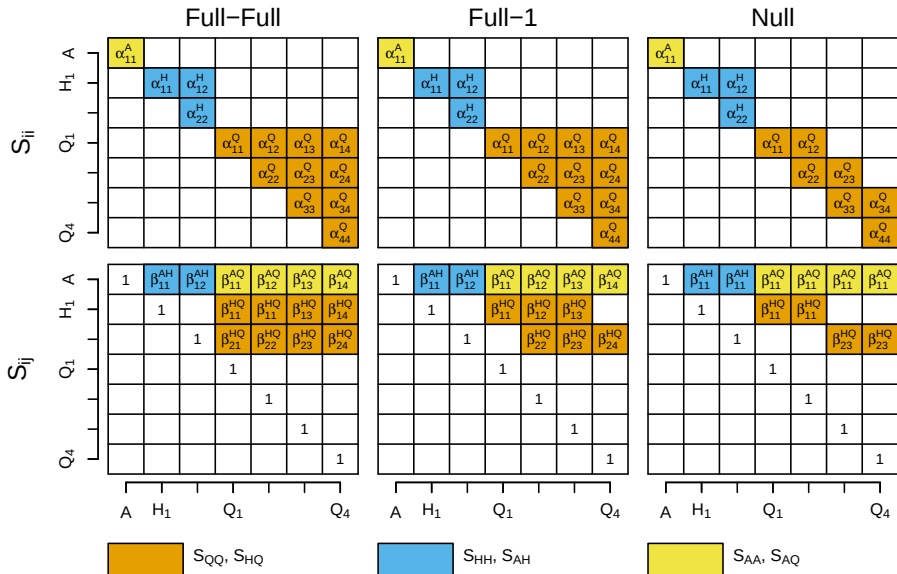
Starting from the bottom level define (Møller et al., 2023)

$$\begin{aligned}
 \boldsymbol{\epsilon}_1 &= \boldsymbol{u}_1 & ; \quad \boldsymbol{u}_1 &\sim N(\mathbf{0}, \boldsymbol{\Sigma}_1), \\
 \boldsymbol{\epsilon}_2 &= \boldsymbol{S}_{21}\boldsymbol{u}_1 + \boldsymbol{u}_2 & ; \quad \boldsymbol{u}_2 &\sim N(\mathbf{0}, \boldsymbol{\Sigma}_2), \\
 \boldsymbol{\epsilon}_3 &= \boldsymbol{S}_{31}\boldsymbol{u}_1 + \boldsymbol{S}_{32}\boldsymbol{u}_2 + \boldsymbol{u}_3 & ; \quad \boldsymbol{u}_3 &\sim N(\mathbf{0}, \boldsymbol{\Sigma}_3), \\
 &\vdots & \\
 \boldsymbol{\epsilon}_K &= \sum_{j=1}^{K-1} \boldsymbol{S}_{Kj}\boldsymbol{u}_j + \boldsymbol{u}_K & ; \quad \boldsymbol{u}_K &\sim N(\mathbf{0}, \boldsymbol{\Sigma}_K),
 \end{aligned}$$

where $Cov[\boldsymbol{u}_i, \boldsymbol{u}_j] = \mathbf{0}$, $(i \neq j)$. And

$$\boldsymbol{\Sigma}_i^{-1} = \boldsymbol{S}_{ii}\boldsymbol{S}_{ii}^T,$$

with $\boldsymbol{S}_{ii}$ is an upper triangular matrix.



Estimation

Likelihood

$$l(\boldsymbol{\beta}, \boldsymbol{\alpha}; \mathbf{V}) \propto -\frac{T-1}{2} \text{Tr} \boldsymbol{\Sigma}^{-1} \mathbf{V} + \frac{T-1}{2} \log |\boldsymbol{\Sigma}^{-1}|,$$

where $\mathbf{V}$ is the observed variance-covariance matrix, $\boldsymbol{\Sigma}^{-1}$ is parameterized through the model and estimation is done

- Sequentially starting from the bottom level
- Using:
 - a robust EM-like algorithm (solving normal equations) and
 - the Newton method (using the Hessian of the likelihood)

diagonal elements of $\mathbf{S}_{ii}$ is estimated in the log-domain.

Shrinkage

A simple modification of the likelihood by introducing weights in the following way:

$$l_s(\boldsymbol{\Sigma}; \mathbf{V}, \mathbf{w}) = l(\boldsymbol{\Sigma}; w_1 \mathbf{V} + w_2 \text{blockdiag} \mathbf{V} + w_3 \text{diag} \mathbf{V}),$$

where $\sum_i w_i = 1$.

Statistical tests

As the method is based on likelihood estimation we have access to

- Wald test for individual parameters or sets of parameters
- Likelihood ratio test for individual parameters or specific hypothesis

In the work we explore

- Wald test for testing if parameters should be zero or equal and confirm using LRT
- Effect of different initial structures

The algorithm, Bottom level

1 Estimation.

- 1 Initialise a variance–covariance structure for the bottom level
- 2 Iterate between α_{ij} and α_{ii} a fixed number of times
- 3 Use the result from 2 as initial value for the Newton method.
- 4 Report the parameter estimates, likelihood, and Hessian.

2 Model reduction.

- 1 Choose in which order to test parameters for removal
- 2 Calculate the test statistics for the increasing set of parameters until significant and confirm by LRT.
- 3 Calculate the test statistics for all relevant pairwise comparison and individual parameters that can be set to zero
- 4 Reduce model, and confirm by likelihood-ratio test.
- 5 Report parameters and structure of S_{ii} and S_{ij} .

The algorithm, Top levels

- 1 Iterate through higher aggregation levels with the lower levels fixed in the same way as the bottom level, now including estimation of β
- 2 Report the final result.

Case study

- Electricity load in Sweden 2016-2020
- 2016-2019 used for estimating mean value structure
- Linear model including annual and diurnal variation
- Double seasonal AR models of the residuals (daily and weekly)
- Temporal reconciliation of residuals



Some results

	SE		SE1		SE2		SE3		SE4	
	df	RRMSE	df	RRMSE	df	RRMSE	df	RRMSE	df	RRMSE
Obs-test	1830	-12.4	1830	-8.5	1830	-9.8	1830	-15.9	1830	-15.3
Obs-train	1830	0.1	1830	4.1	1830	2.1	1830	-2.3	1830	-5.2
Full-Full ₁	402	-4.8	306	-1.7	413	-4.1	413	-5.6	464	-5.5
Full-Full ₂	402	-5.3	306	0.8	413	-3.1	413	-5.8	464	-6.6
Full-2 ₁	219	-3.8	180	-3.2	237	-4.0	240	-3.9	264	-4.0
Full-2 ₂	219	-3.8	180	-2.2	237	-3.8	240	-4.9	264	-3.7
Null-Null ₁	83	-4.3	77	-2.7	81	-4.1	85	-4.8	85	-4.6
Null-Null ₂	83	-3.7	77	-0.7	81	-3.1	85	-4.0	85	-4.1
Shrink	0.015	-5.3	0.016	-1.3	0.035	-4.0	0.015	-7.4	0.016	-7.0
Auto-cov.	69	-2.2	49	-2.9	49	-3.5	71	-1.6	71	-2.2
Model-AR1	60	-2.4	60	-3.0	60	-3.2	60	-1.9	60	-2.2
Diag	60	-3.1	60	-2.4	60	-2.6	60	-2.9	60	-2.6

Outline

- 1 Introduction
- 2 Heat load forecasting
- 3 Likelihood inference
- 4 The general linear model**
- 5 Conclusion

A general linear model formulation

Consider the general linear model

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}; \quad \boldsymbol{\epsilon} \sim N(0, \mathbf{I}_T \otimes \boldsymbol{\Sigma}_r).$$

where $\mathbf{y}^T = [\mathbf{y}_T^T, \mathbf{y}_{T-1}^T \cdots, \mathbf{y}_1^T]^T$.

Question: Can we find a formulation of $\mathbf{X}$ such that $\boldsymbol{\beta}$ and

$$\mathbf{P} = (\mathbf{S}^T \boldsymbol{\Sigma}^{-1} \mathbf{S})^{-1} \mathbf{S}^T \boldsymbol{\Sigma}^{-1}$$

are equivalent?

Theorem (Equivalence with hierarchical reconciliation)

The model,

$$\mathbf{y}_t = \mathbf{X}_{\cdot,t} \boldsymbol{\beta} + \boldsymbol{\epsilon}_t; \quad \boldsymbol{\epsilon}_t \sim N(0, \boldsymbol{\Sigma}_r),$$

with $\mathbf{X}_{\cdot,t} = \mathbf{I} \otimes \hat{\mathbf{y}}_t$, and linear constraints $\text{vec}^{-1}(\boldsymbol{\beta})^T \mathbf{S} = \mathbf{I}$, is equivalent to the hierarchical reconciliation, formally

$$\mathbf{P} = \text{vec}^{-1}(\hat{\boldsymbol{\beta}})^T,$$

where $\mathbf{P}$ is the usual hierarchical reconciliation with $\boldsymbol{\Sigma}_h = \frac{1}{T}(\mathbf{Y} \mathbf{S}^T - \hat{\mathbf{Y}})^T(\mathbf{Y} \mathbf{S}^T - \hat{\mathbf{Y}})$. The constraint problem can further be formulated as

$$\mathbf{y}_t - \hat{\mathbf{y}}_{bot,t} = [\mathbf{I}_m \otimes (\hat{\mathbf{y}}_{T,t}^T - \hat{\mathbf{y}}_{bot,t}^T \mathbf{S}_T^T)] \boldsymbol{\beta}_T + \boldsymbol{\epsilon}_t; \quad \boldsymbol{\epsilon}_t \sim N(\mathbf{0}, \boldsymbol{\Sigma}_r).$$

Comments

Here

$$\hat{\mathbf{y}} = \begin{bmatrix} \hat{\mathbf{y}}_{\text{top},t} \\ \hat{\mathbf{y}}_{\text{bot},t} \end{bmatrix}, \quad \mathbf{S} = \begin{bmatrix} \mathbf{S}_{\text{top}} \\ \mathbf{I}_m \end{bmatrix}.$$

The full model can be written as

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}; \quad \boldsymbol{\epsilon} \sim N(\mathbf{0}, \mathbf{I}_T \otimes \boldsymbol{\Sigma}_r)$$

and the MLE of $\boldsymbol{\beta}$ is independent of $\boldsymbol{\Sigma}_r$, and

$$\hat{\boldsymbol{\Sigma}}_{r,\text{REML}} = \frac{1}{T - (n - m)} \sum_{t=1}^T \mathbf{e}_t \mathbf{e}_t^\top,$$

where $\mathbf{e}_t = \mathbf{y}_t - \mathbf{X}_{:,t} \hat{\boldsymbol{\beta}}$.

More comments

Notice that

- 1 Estimation of the $n \times n$ ($\frac{n(n+1)}{2}$ parameters) variance covariance matrix is replaced by estimation of $m(n - m)$ mean value parameters.
- 2 The result is independent of the forecast variance Σ_r , which may be parameterized by some model (e.g. AR(1) structures)

Shrinkage

Usually a shrinkage estimator for Σ is used, in order to improve the condition number of Σ .

- In the regression setting an ill conditioned covariance matrix will show up as multicollinearity.
- Strong multicollinearity is usually handled by regularizations (e.g. Ridge regression) or equivalent by introducing parameter priors.
- Specific choices of priors (depending on the shrinkage) lead to equivalence between the regression setting and the usual (called MinT) approach.

Theorem

The weight matrix $\mathbf{P}(\lambda)$ is equivalent to the linear regression model

$$\mathbf{y}_t - \hat{\mathbf{y}}_{B,t} = [\mathbf{I}_m \otimes (\hat{\mathbf{Y}}_{T,t} - \hat{\mathbf{Y}}_{B,t} \mathbf{S}_T^T)] \boldsymbol{\beta}_T + \boldsymbol{\epsilon}_t; \quad \boldsymbol{\epsilon}_t \sim N(\mathbf{0}, \boldsymbol{\Sigma}_r),$$

with prior distribution $\boldsymbol{\beta}_T \sim N(\boldsymbol{\beta}_{0,T}, \boldsymbol{\Sigma}_r \otimes \boldsymbol{\Sigma}_{\beta,0})$, and

$$\boldsymbol{\beta}_{0,T} = \text{vec} \left[\left(\boldsymbol{\Sigma}_{h,T}^d + \mathbf{S}_T \boldsymbol{\Sigma}_{h,B}^d \mathbf{S}_T^T \right)^{-1} \mathbf{S}_T \boldsymbol{\Sigma}_{h,B}^d \right]$$

$$\boldsymbol{\Sigma}_{0,\beta} = \frac{1 - \lambda}{\lambda T} (\boldsymbol{\Sigma}_{h,T}^d + \mathbf{S}_T \boldsymbol{\Sigma}_{h,B}^d \mathbf{S}_T^T)^{-1},$$

Residual variance–covariance

To obtain complete equivalence priors for the residual variance Σ_r is also needed (related to the Inverse Wishart distribution).

Using the regression formulation (including priors), we arrive at two variance estimators, one with a REML correction term and one without

$$\hat{\Sigma}_{r,\text{shrink}} = \mathbf{P}(\lambda) \hat{\Sigma}_h \mathbf{P}^\top(\lambda)$$

$$\hat{\Sigma}_{r,\text{sREML}} = \frac{T}{T - (n - m)(1 - \lambda)} \mathbf{P}(\lambda) \hat{\Sigma}_s \mathbf{P}^\top(\lambda).$$

Forecast variance

In general the forecast variance can be written as

$$\hat{V}[\mathbf{y}_{t+h} - \tilde{\mathbf{y}}_{t+h}] = \hat{V}[\mathbf{X}_{\cdot, t+h} \hat{\boldsymbol{\beta}} + \boldsymbol{\epsilon}_{t+h}] = \mathbf{X}_{\cdot, t+h} \hat{V}[\hat{\boldsymbol{\beta}}] \mathbf{X}_{\cdot, t+h}^{\top} + \hat{\boldsymbol{\Sigma}}_r.$$

with

$$\hat{V}[\hat{\boldsymbol{\beta}}_{\text{top}}] = \hat{\boldsymbol{\Sigma}}_r \otimes \left((\mathbf{X}_{1,\cdot}^{\top} \mathbf{X}_{1,\cdot} + \boldsymbol{\Sigma}_{\beta,0})^{-1} \mathbf{X}_{1,\cdot}^{\top} \mathbf{X}_{1,\cdot} (\mathbf{X}_{1,\cdot}^{\top} \mathbf{X}_{1,\cdot} + \boldsymbol{\Sigma}_{\beta,0})^{-1} \right).$$

The parameters (weights) uncertainty and the parameters correlation can be assessed from the using the above, this may be used for testing and further uncertainty may propagate through to predictions variances.

Simulation study

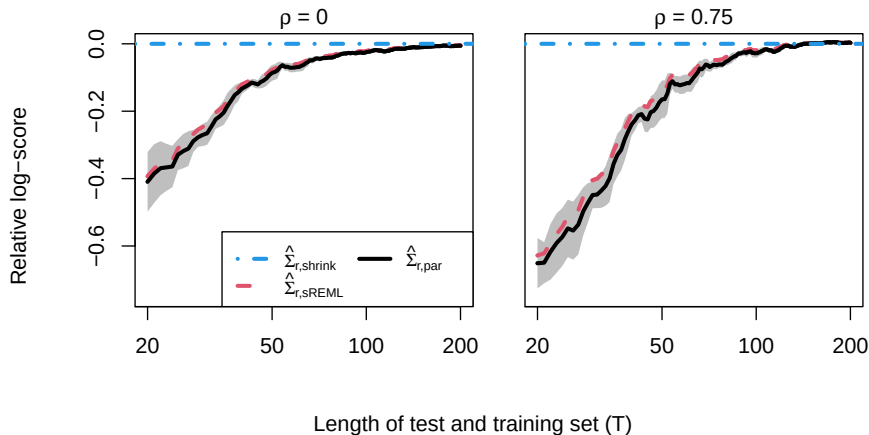
We simulate the data-generating process at the bottom level ($\mathbf{y}_t \in \mathbb{R}^4$) using a multivariate AR(1) process

$$\mathbf{y}_t = \mathbf{A}\mathbf{y}_{t-1} + \boldsymbol{\epsilon}_t; \quad \boldsymbol{\epsilon}_t \sim N(\mathbf{0}, \boldsymbol{\Sigma}_\epsilon) \quad \text{and iid.}, \quad (2)$$

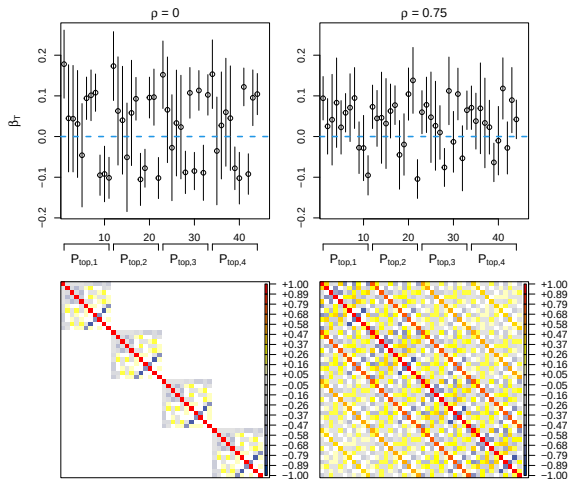
with $\mathbf{A}_{ii} = 0.7$, $\mathbf{A}_{ij} = -0.25$ ($i \neq j$), and $(\boldsymbol{\Sigma}_\epsilon)_{ij} = \rho^{|i-j|}$, here we present simulation studies using $\rho = 0$ and $\rho = 0.75$.

Univariate AR(1) models for all levels in the hierarchy (all combinations of summations of 1, 2, 3, or 4 variables).

Simulation results



Parameter standard error and correlation¹



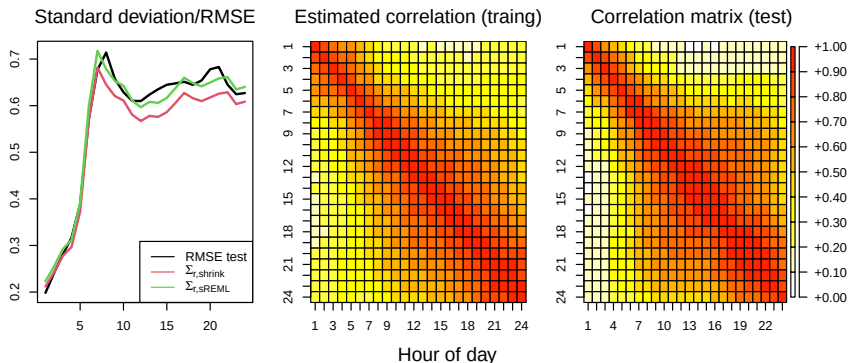
¹moller2024

Case study

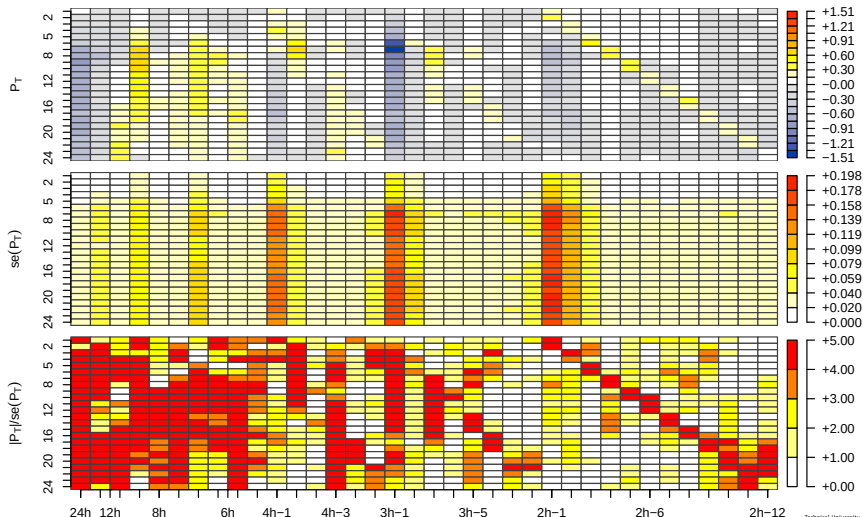
- Electricity load in Sweden 2016-2020
- Linear model including annual and diurnal variation
- Double seasonal AR models of the residuals (daily and weekly)
- Temporal reconciliation of residuals



Case: variance–covariance



Case: parameters



Outline

- 1 Introduction
- 2 Heat load forecasting
- 3 Likelihood inference
- 4 The general linear model
- 5 Conclusion**

Conclusion / Future

There are many articles on the development of better estimators coming out, and all is not finalized as of yet. Some points are

- Reconciliation ensures coherent forecasts on all levels
- Reconciliation very often improves forecast skills
- The covariance matrix is often ill-conditioned and hence shrinkage should be applied (in the regression setting this is multicollinearity)

My personal interest for future development is

- Parametrized (time varying) residual variance covariance matrices
- Include auto-correlation between time steps
- Develop test strategies to reduce model complexity

References

- [1] Bergsteinsson H.G., Møller J.K., Nystrup P., Pálsson Ó.P., Guericke D., and Madsen H. (2021). Heat load forecasting using adaptive temporal hierarchies. *Applied Energy*, **292**.
- [2] Møller J.K., Nystrup P., Hjort P.G., and Madsen H. (2024). Optimal Forecast Reconciliation with Uncertainty Quantification [arXiv:2402.06480](https://arxiv.org/abs/2402.06480).
- [3] Møller J.K., Nystrup P., and Madsen H. (2023). Likelihood-based inference in temporal hierarchies. *International Journal of Forecasting*, *in press*.
- [4] Møller JK, Nystrup P, and Madsen H. (2022) Supplementary material for "Likelihood-based inference in temporal hierarchies". <http://people.compute.dtu.dk/jkmo/>
- [5] Nystrup, P., Lindström, E., Møller, J. K., and Madsen, H. (2021). Dimensionality reduction in forecasting with temporal hierarchies. *International Journal of Forecasting*, *37*(3)
- [6] Nystrup, P., Lindström, E., Pinson, P., and Madsen, H. (2020). Temporal hierarchies with autocorrelation for load forecasting. *European Journal of Operational Research*, *280*, 876-888.