

Short-term forecast of the geomagnetic secular variation using recurrent neural networks trained by Kalman filter

Sho Sato , Hiroaki Toh

Division of Earth and Planetary Sciences, Graduate School of Science,
Kyoto University, JAPAN

Wed, 18 Sep. 09:20-09:40 | IUGONET 研究集会 @ ハイブリッド

OVERVIEW

1. INTRODUCTION

- Background : *Geomagnetic Secular Variation* and *Geomagnetic Jerk*
- Previous Research : Data Assimilation = Geodynamo SIM x Observed Data

2. METHODOLOGY

- Basics of Machine Learning and *Recurrent Neural Networks*
- Overfitting problem in the Backpropagation algorithm
- The extended Kalman filtering approach to train *RNNs*

3. RESULTS SO FAR

- Hindcast (training: 2004.75 – 2014.50 → validation: 2014.75 – 2019.50)
- Forecast (training: 2014.37 – 2024.13 → validation: 2024.37 – 2029.13)

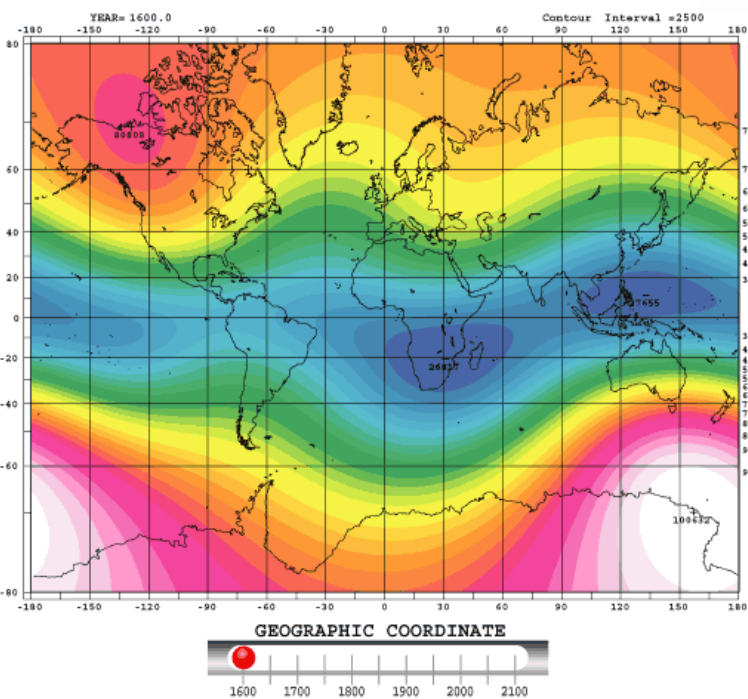
4. RESEARCH PLAN

- Conclusion & Road Map

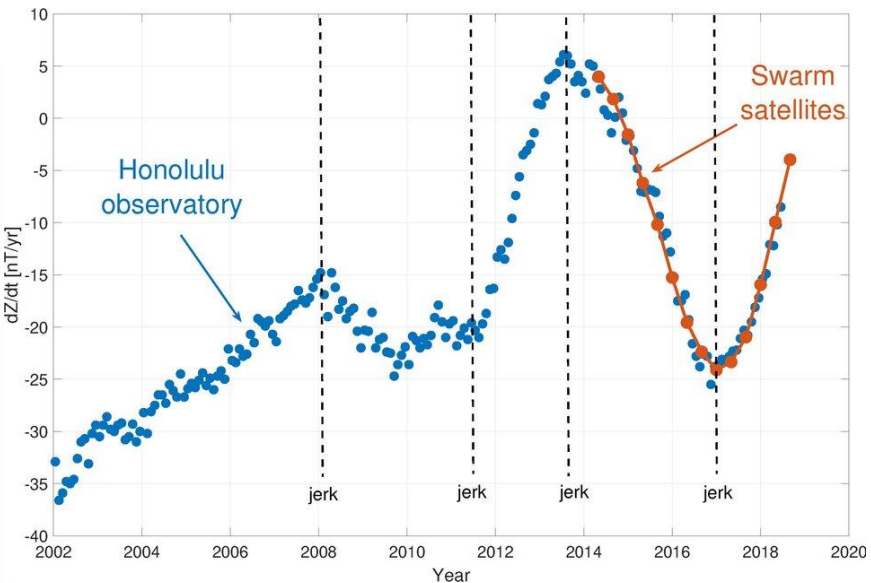
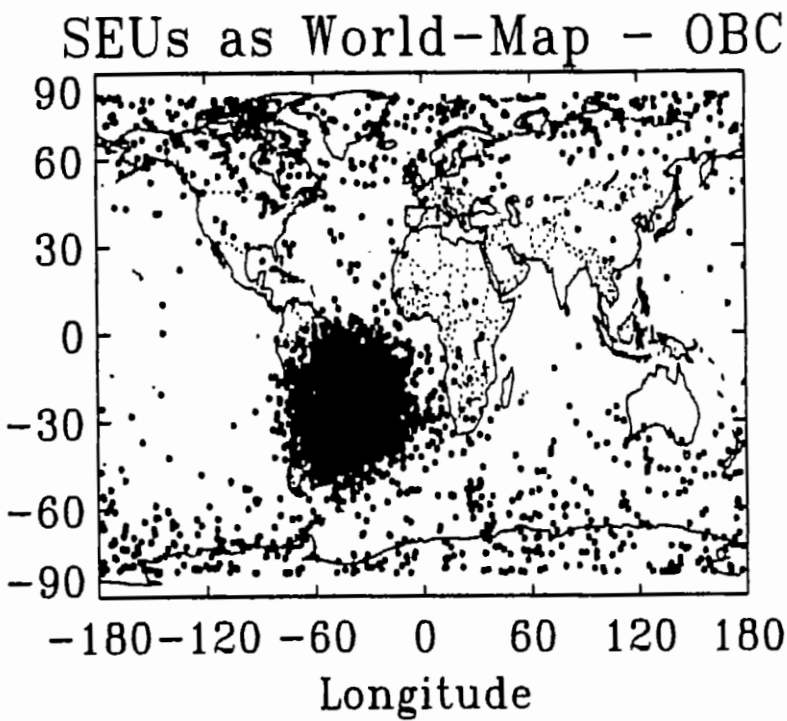
INTRODUCTION | Background

Geomagnetic *Secular Variation*

(Image: WDC for Geomagnetism, Kyoto*¹)



Location of UOSat-2 when *Single-event Upsets* were detected in its OBC memory during 1989
(Image: Daly et al, 1996*²)



Geomagnetic *Jerks* (= Sudden Acceleration)
(Image: ESA/DTU*³)

The *Gauss coefficients* and their *derivatives* have been reported every 5 years as *International Geomagnetic Reference Field*.

**Gauss coefficients* ($g_0^1, g_1^1, h_1^1, \dots$) = Spherical harmonic coefficients of the Geomagnetic scalar potential, equivalent to multipole moments.

1. <https://wdc.kugi.kyoto-u.ac.jp>, 2. Daly et al (1996), "Space Environment Analysis: Experience and Trends" vol.392. p.15., 3. [Swarm helps explain Earth's magnetic jerks](#).

INTRODUCTION | Previous Research

IGRF-13 candidate SV models= **Data Assimilation** (**Geodynamo SIM** x Observation)

	Forecast Error (RMSE [nT])	Hindcast	Simulation Model	Observational Data
Sanchez et al. (2020)* ¹	50 ~ 100	1840 ~ 2020	3D dynamo x 100 ens.	COV-OBS.x1 Kalmag model
Minami et al. (2020)* ²	100.9 ~ 228.5	2004 ~ 2019	MHD dynamo x 1000 ens.	MCM model
Tangborn et al. (2021)* ³	$O(10^{1.5} \sim 10^2)$	1590 ~ 2015	NASA GEMS x 400 ens.	gufm1, CM5, CM6

MOTIVATION : Can **ML** models be a new approach to forecast SV without simulations?

*1. doi: [10.1186/s40623-020-01279-y](https://doi.org/10.1186/s40623-020-01279-y) , *2. doi: [10.1186/s40623-020-01253-8](https://doi.org/10.1186/s40623-020-01253-8) , *3. doi: [10.1186/s40623-020-01324-w](https://doi.org/10.1186/s40623-020-01324-w).

INTRODUCTION | ML / DL in Geophysics

Reviews of Geophysics

REVIEW ARTICLE
10.1029/2021RG000742

Deep Learning for Geophysics: Current and Future Trends

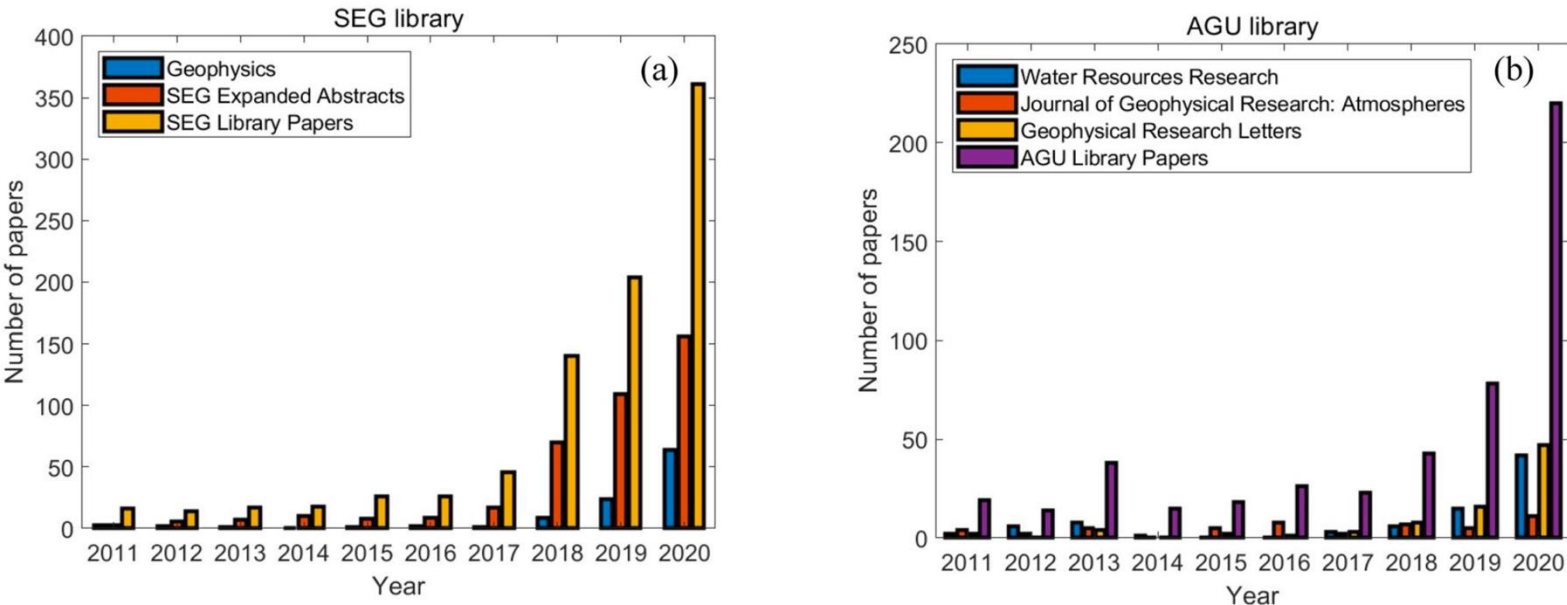


Figure 3. (a) and (b) are statics of artificial intelligence (AI)-related papers in SEG Library and AGU Library. In (a), Geophysics means the flagship journal of SEG. SEG Expanded Abstracts means the Expanded Abstracts from SEG annual meeting. SEG Library papers mean the papers founded in the SEG digital library. In (b), the first three captions in the legend are the names of top journals in AGU. The fourth caption in the legend represents the papers founded in the AGU digital library.

“The geophysical community has shown great interests in DL in recent years.

Figure 3 shows the published papers related to artificial intelligence in two major geophysical unions, that is, society of exploration geophysics (SEG) and American geophysical union (AGU). A clear exponential growth is observed in both libraries due to the use of DL techniques. Moreover, DL has also provided several astonishing results to the geophysical community.”

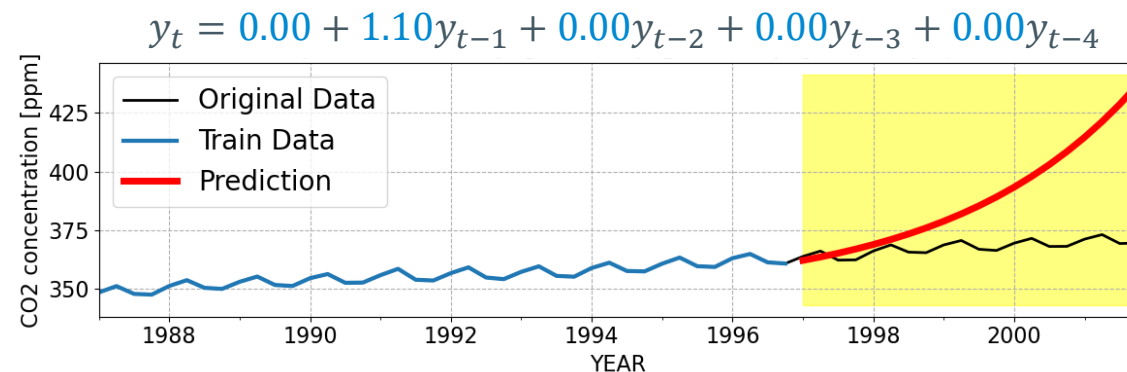
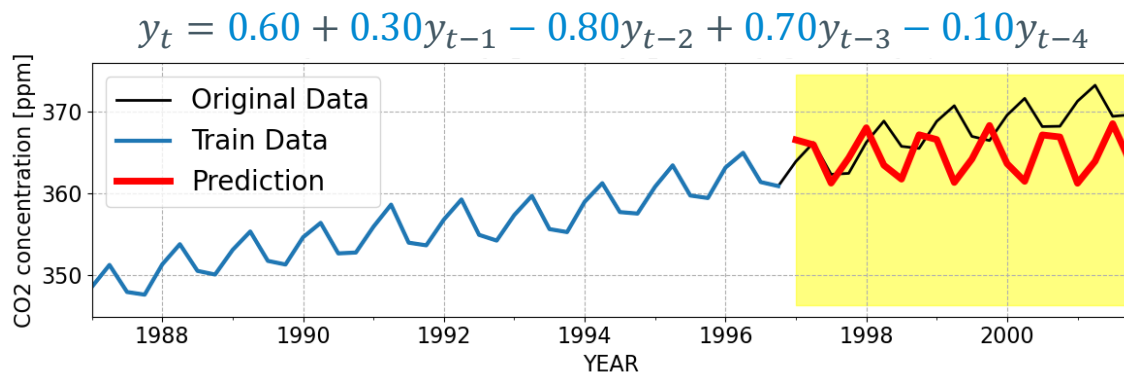
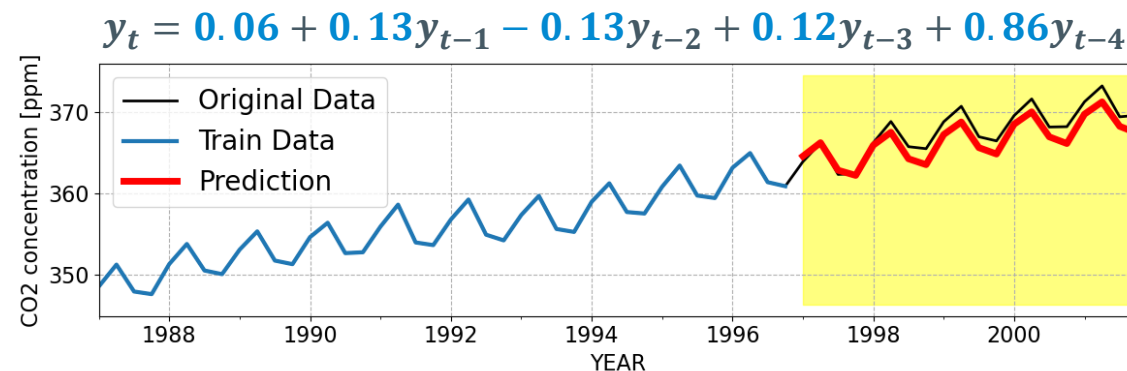
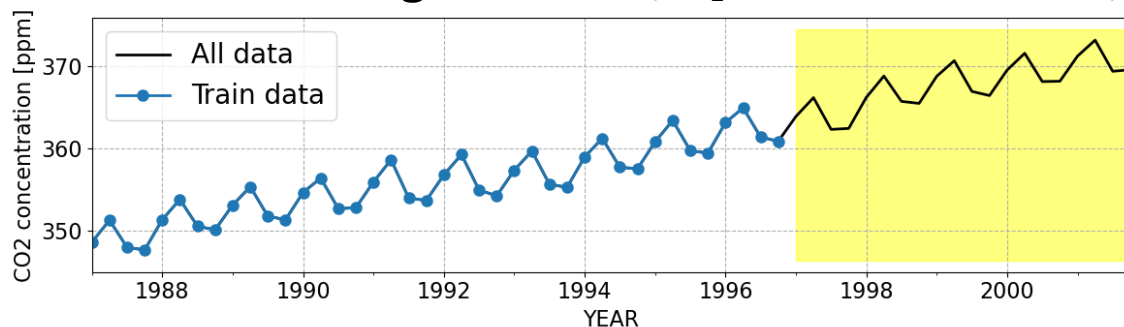
METHOD | What is ML training? = Parameter estimation

Example: AR(4) model

$$y_t = b + w_1 y_{t-1} + w_2 y_{t-2} + w_3 y_{t-3} + w_4 y_{t-4}$$

(Trainable Parameter: b (bias), $w_1 \sim w_4$ (weight))

Original Data (CO₂ at Mauna Loa, Hawaii)



METHOD | Why RNN?

Model		AR	VAR	BPTT-RNNs	Minami et al. (2020) (4D-EnVar)
Average $\sqrt{dP}$		34.48 nT	39.58 nT	<u>33.43 nT</u>	---
5-year $\sqrt{dP}$		88.05 nT	98.33 nT	<u>85.83</u> – 137.74 nT	100.9 – 228.5 nT
Remarks		---	---	$\sqrt{dP}$ depends on the initial $\mathbf{w}_0$ and the # of itr. of BPTT	MHD dynamo(x1000) Depends on the assim. window

METHOD | RNN architecture

Elman Network

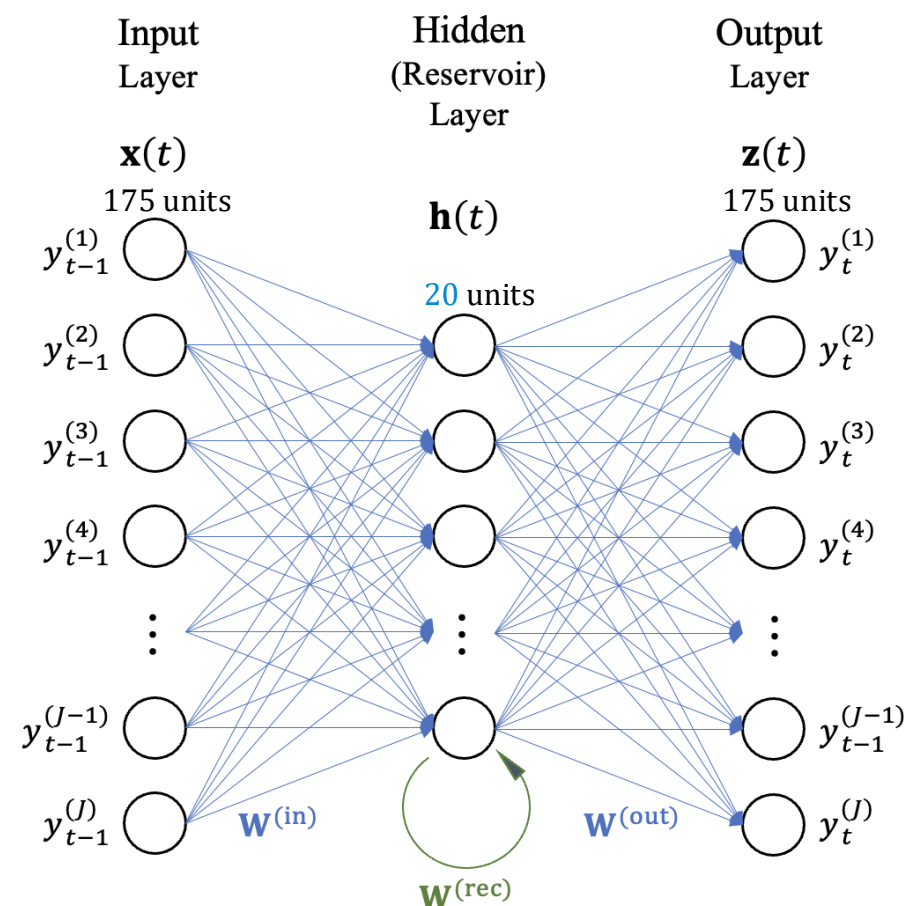
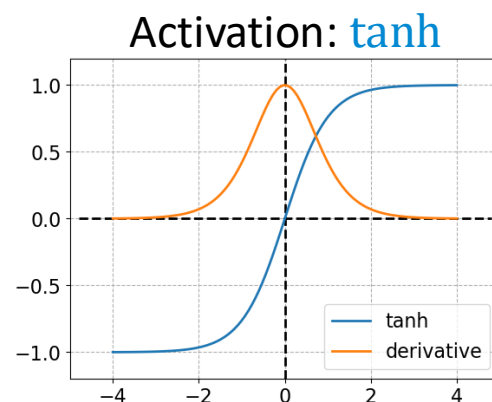
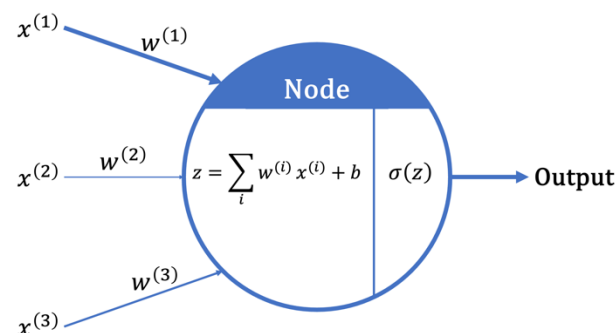
RNN cell

$$\mathbf{h}(t_i) = \tanh(\mathbf{W}^{(\text{in})} \cdot \mathbf{x}(t_i) + \mathbf{W}^{(\text{rec})} \cdot \mathbf{h}(t_{i-1}) + \mathbf{b}^{(\text{rec})})$$

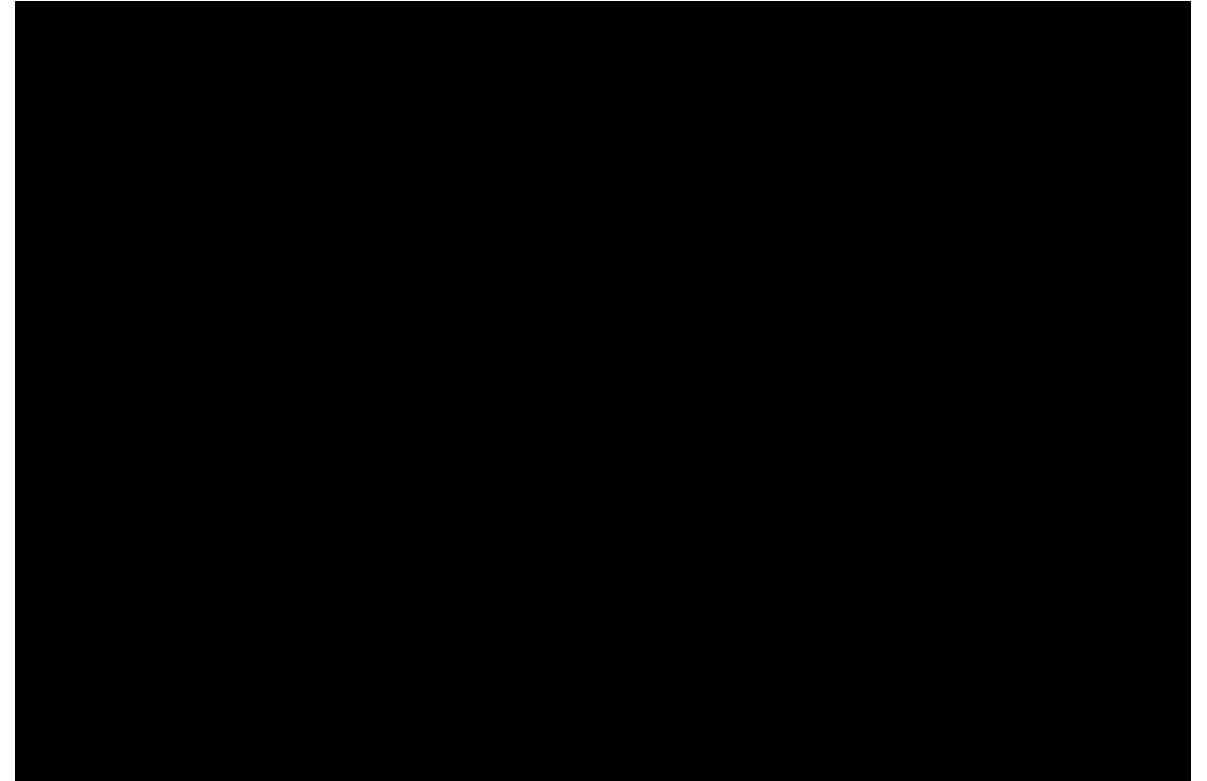
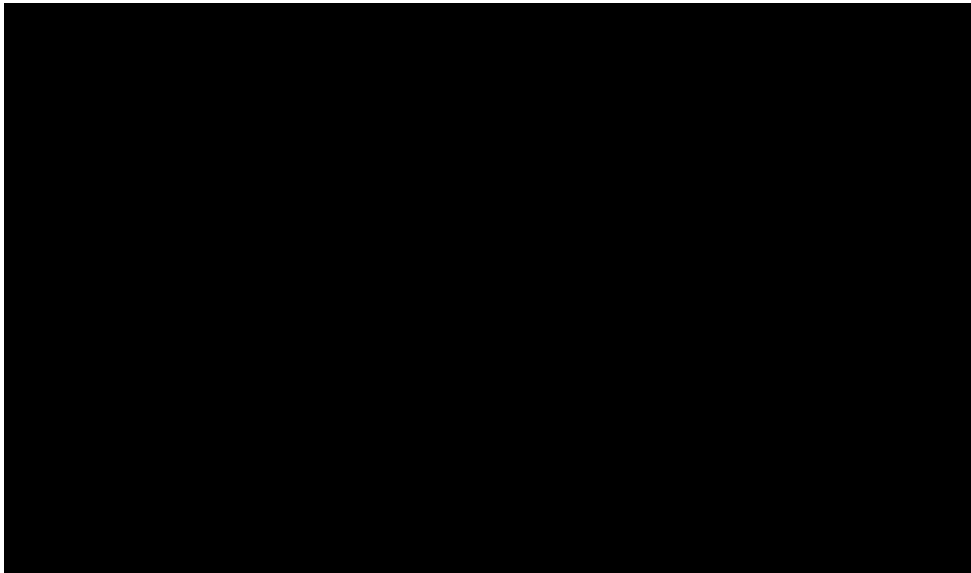
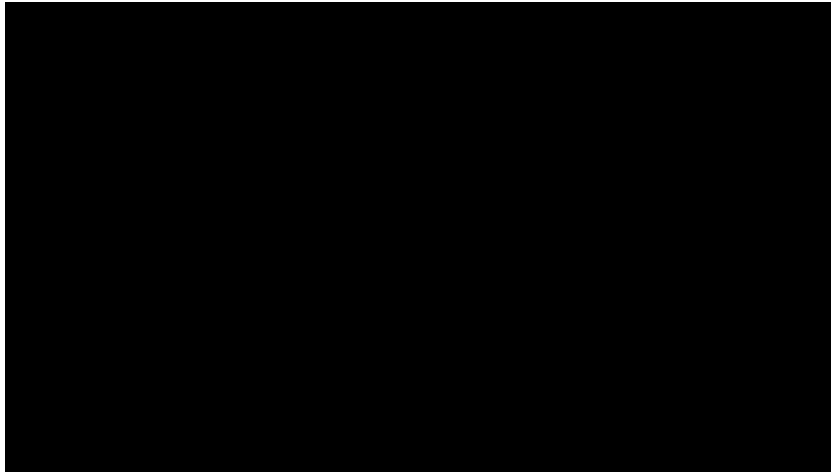
Affine (Dense)

$$\mathbf{z}(t_i) = \mathbf{W}^{(\text{out})} \cdot \mathbf{h}(t_i) + \mathbf{b}^{(\text{out})}$$

(input: $\mathbf{x}(t_i) = [\ddot{g}_n^m(t_{i-1})] \rightarrow$ output: $\mathbf{z}(t_i) = [\ddot{g}_n^m(t_i)]$).



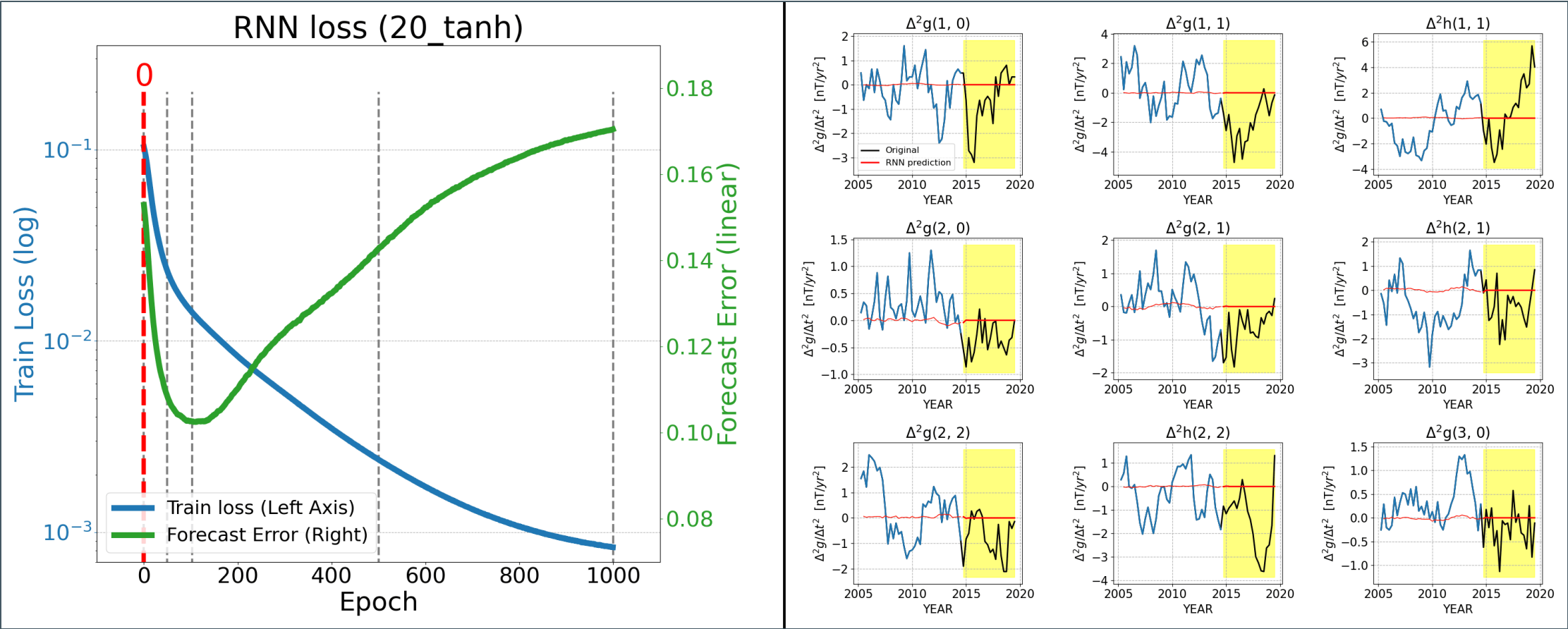
METHOD | Backpropagation algorithm



“What is backpropagation really doing? | Chapter 3, Deep learning” by 3Blue1Brown. (<https://youtu.be/llg3gGewQ5U?si=SUjGzTFV7ELBlo1v>)

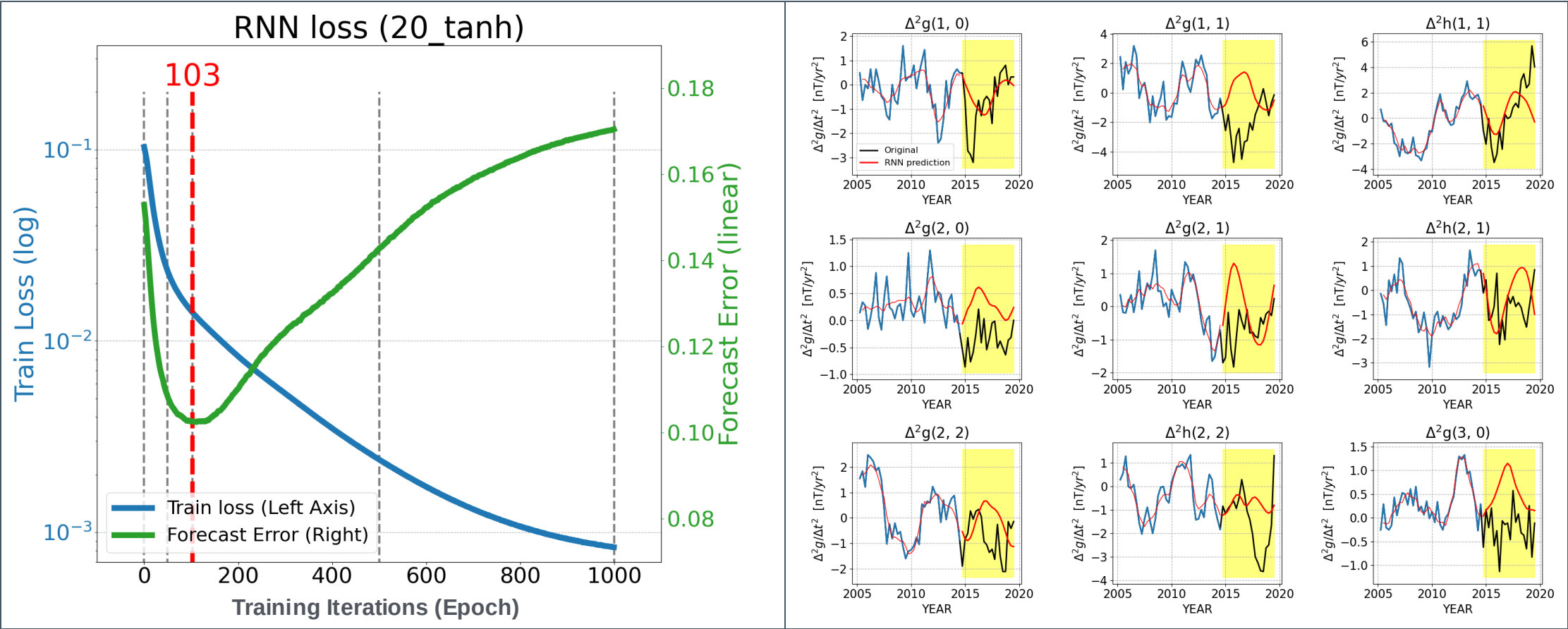
METHOD | Overfitting

Overfitting caused by too much iteration of the Backpropagation algorithm



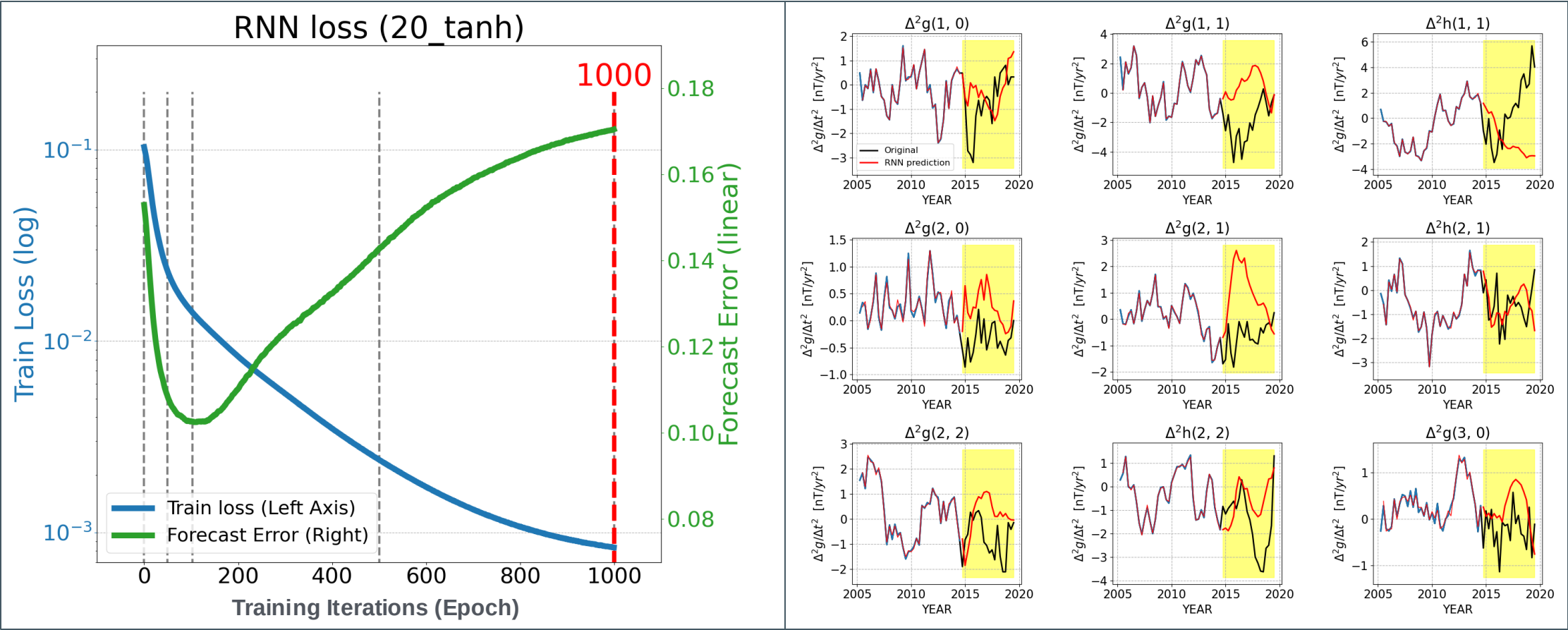
~APR. PROBLEM | Overfitting in RNNs

Overfitting caused by too much iteration of the Backpropagation algorithm



~APR. PROBLEM | Overfitting in RNNs

Overfitting caused by too much iteration of the Backpropagation algorithm



METHOD | EKF algorithm

- State space: $\mathbf{w} = [\mathbf{w}^{(\text{in})}, \mathbf{w}^{(\text{rec})}, \mathbf{b}^{(\text{rec})}, \mathbf{w}^{(\text{out})}, \mathbf{b}^{(\text{out})}]^T$.
- Transition model: $\mathbf{M} = \mathbf{I} \Leftrightarrow \mathbf{w}^f(t) = \mathbf{w}^a(t-1), \mathbf{P}^f(t) = \mathbf{P}^a(t-1)$.

- Observational data: Training data

$\mathbf{y}^o = [g_1^0, g_1^1, h_1^1, \dots, g_{13}^{13}, h_{13}^{13}]^T$, with Observation error - $\mathbf{R}$

- Observation Operator: Forward propagation of RNN

$$\mathcal{H}_t(\mathbf{w}^f(t)) = \text{RNN}(\mathbf{y}^o(t-1) | \mathbf{w}^a(t-1)) = \mathbf{z}(t).$$

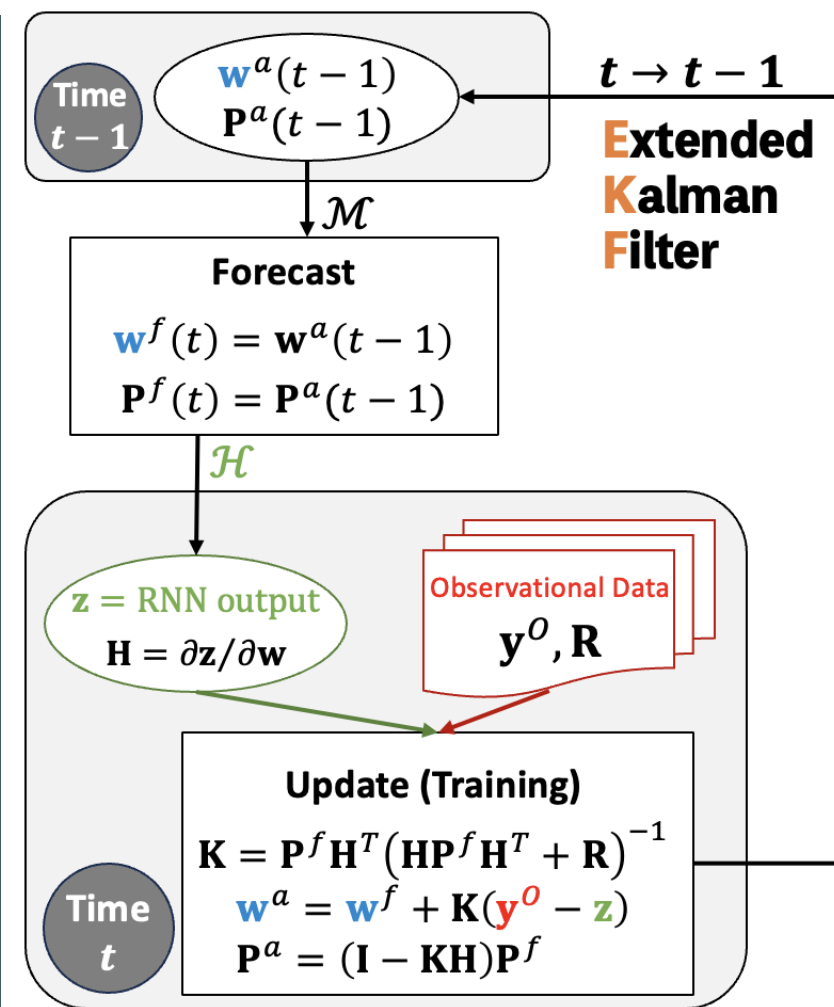
$$\mathbf{H}(t) = \left. \frac{\partial \mathbf{z}}{\partial \mathbf{w}} \right|_{\mathbf{w}=\mathbf{w}^f(t)} = \left[\frac{\partial \mathbf{z}}{\partial \mathbf{W}^{(\text{in})}}, \frac{\partial \mathbf{z}}{\partial \mathbf{W}^{(\text{rec})}}, \dots, \frac{\partial \mathbf{z}}{\partial \mathbf{b}^{(\text{out})}} \right].$$

- Update equation: Weight update (= Training)

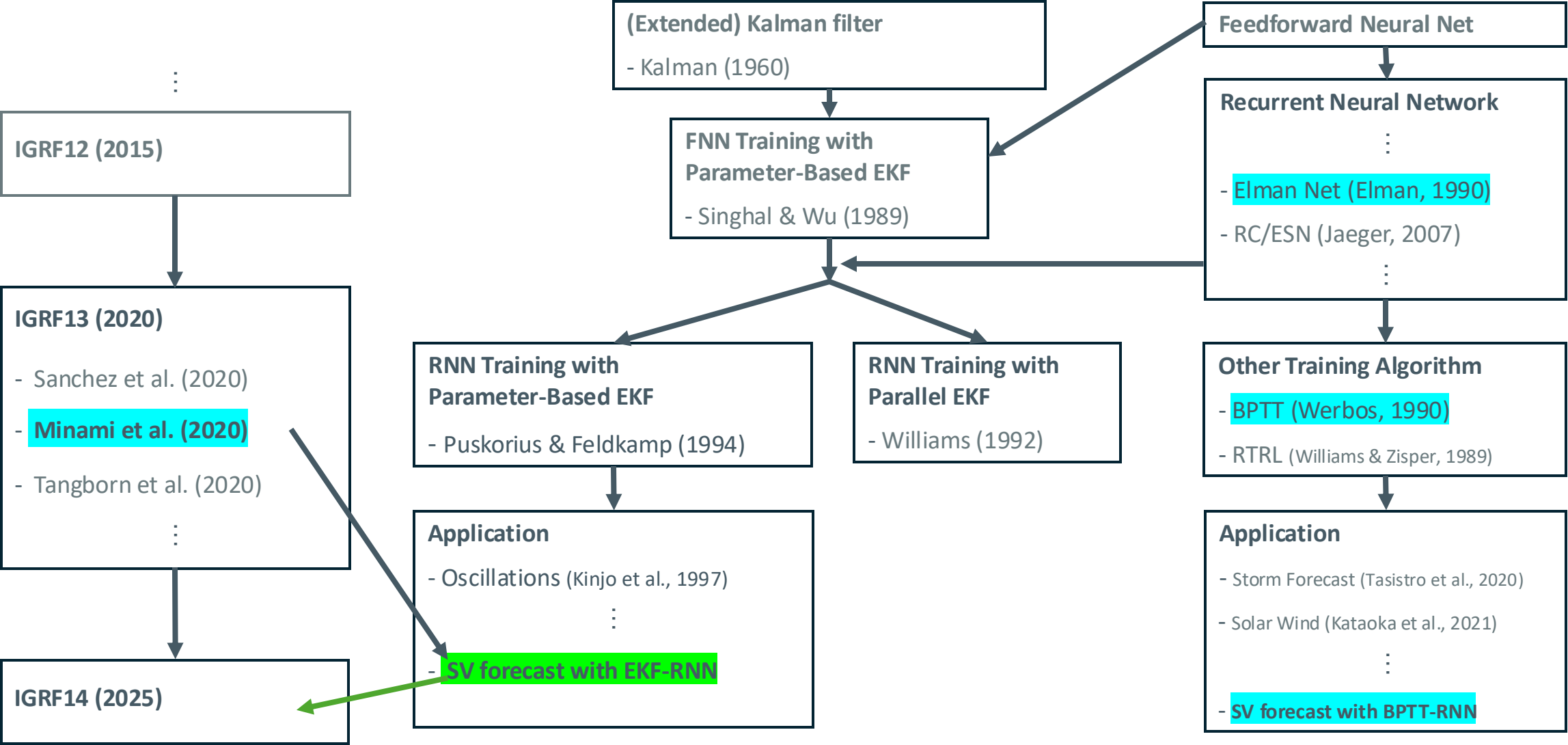
$$\mathbf{w}^a(t) = \mathbf{w}^a(t-1) + \mathbf{K}(t)(\mathbf{y}^o(t) - \mathbf{z}(t)), \quad \mathbf{P}^a(t) = (\mathbf{I} - \mathbf{K}(t)\mathbf{H}(t))\mathbf{P}^f(t).$$



Point : ① Forecast error - $\mathbf{P}$. ② Prevents Overfitting.



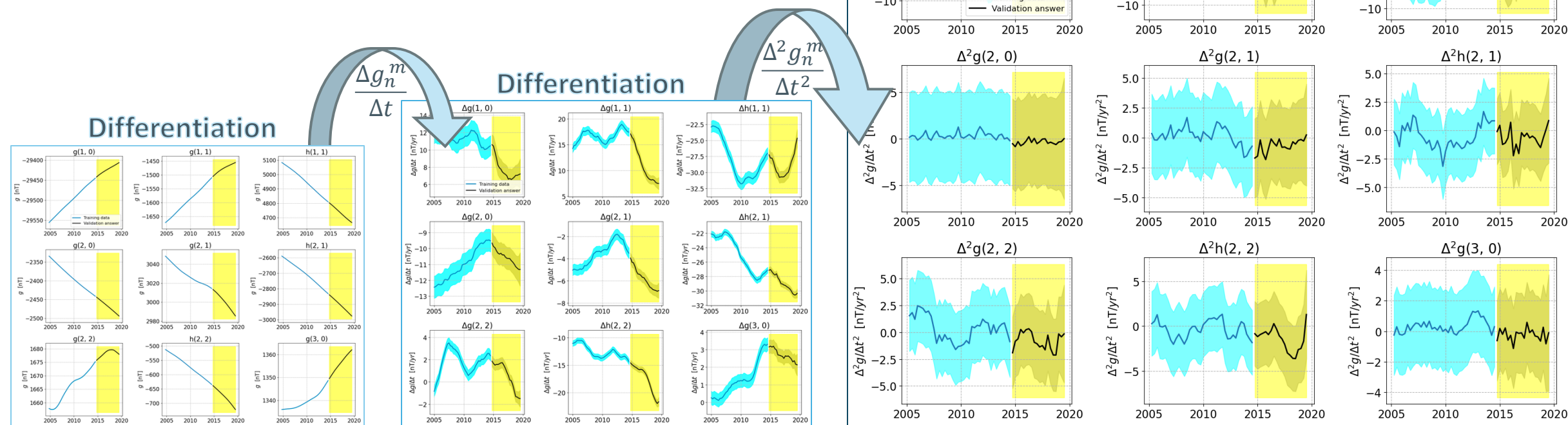
METHOD | Research Context



METHOD | Hindcast w/ MCM-2020

We assess **RNNs** under the same evaluation setup as in Minami et al. (2020)

Input Data : MCM Model (Ropp et al., 2020)
 Training Period : 2004.75 ~ 2014.50 (40 points)
 Validation Period : 2014.75 ~ 2019.50 (20 points)



HINDCAST | 2nd order time derivative ($\ddot{y}$)



Common Feature:

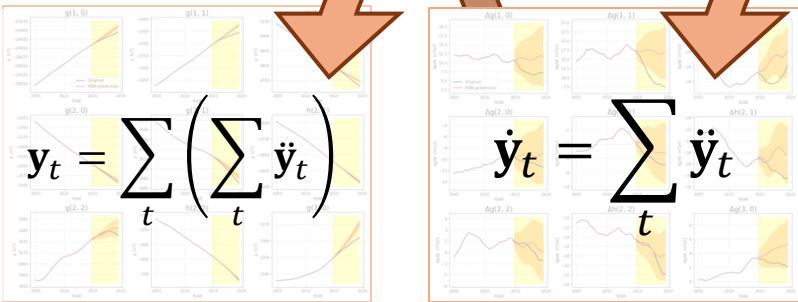
Predicted dynamics = **sin-like** wave



Inaccurate Component (Ex. g_1^0, g_3^0):

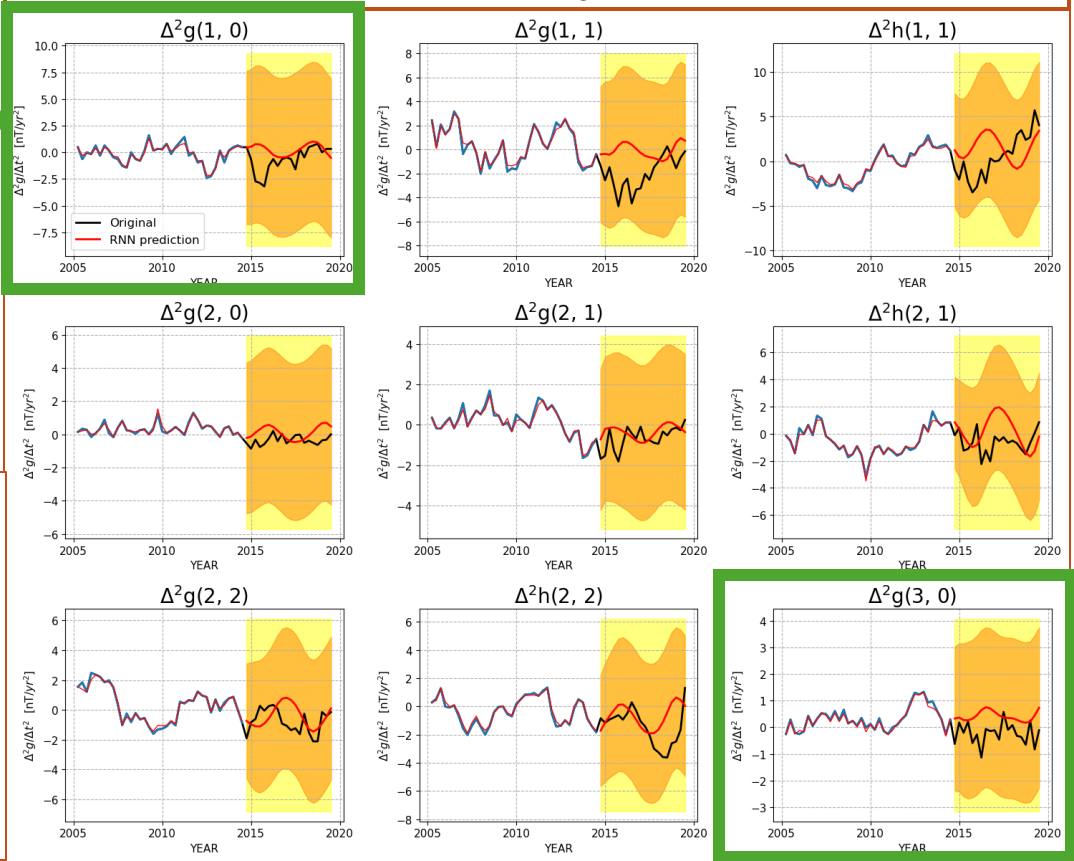
Incorrect positive and negative sign

Summation



- Training Data (MCM model)
- Validation Data (MCM model)
- RNN Prediction

$$\ddot{y} = \frac{\Delta^2 g_n^m}{\Delta t^2}$$



HINDCAST | 1st order time derivative ($\dot{y}$)



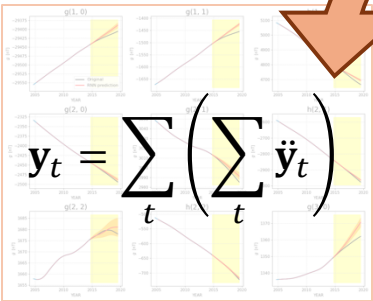
Common Feature:
~~Extends the last trend~~



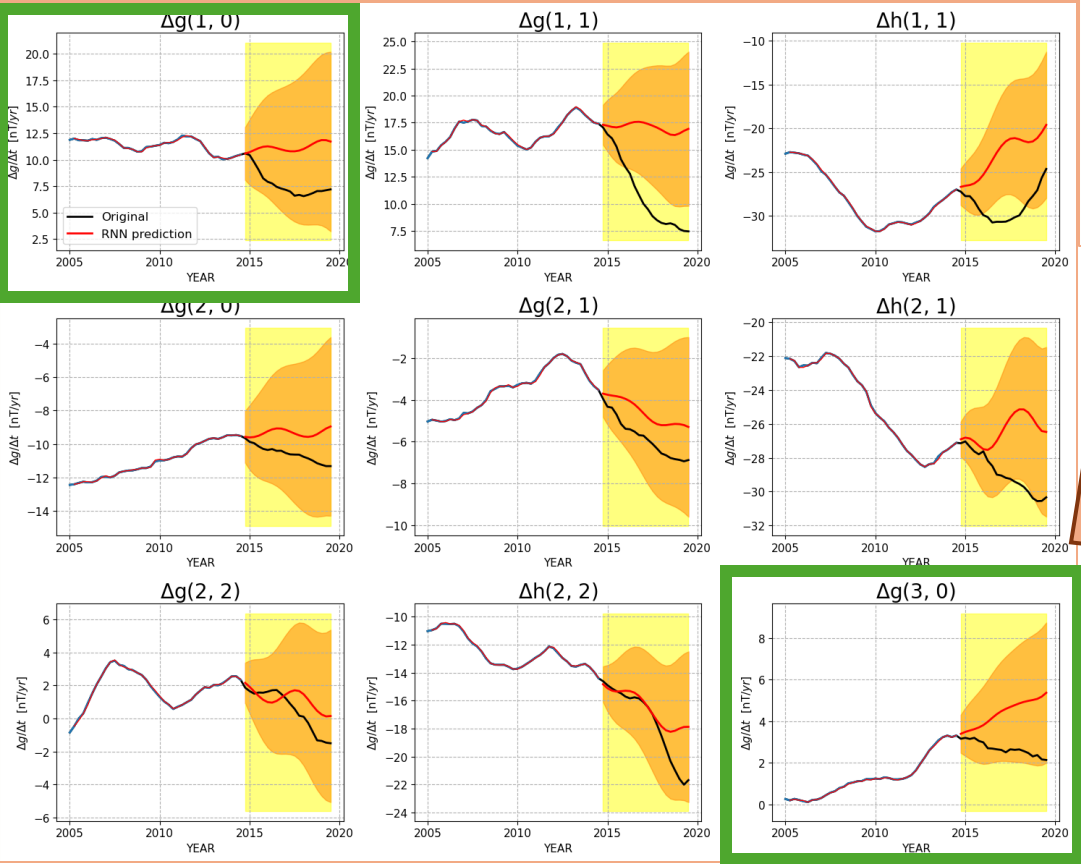
(Ex. g_2^0, g_3^0)
Inaccurate Component:

Inaccurate forecast caused by
the sudden change of the trend
in the validation period (2014).

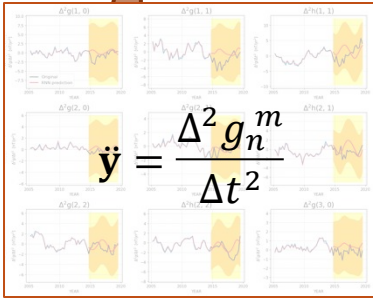
Summation



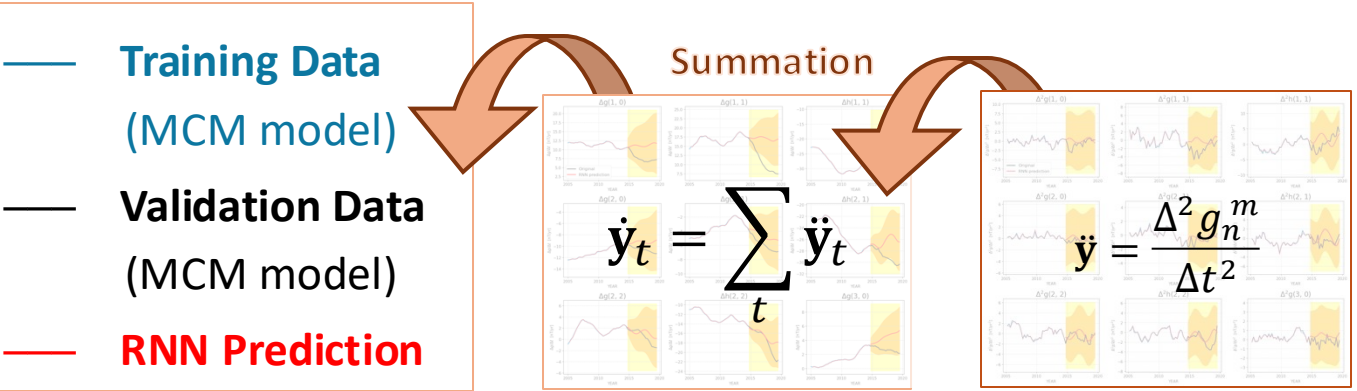
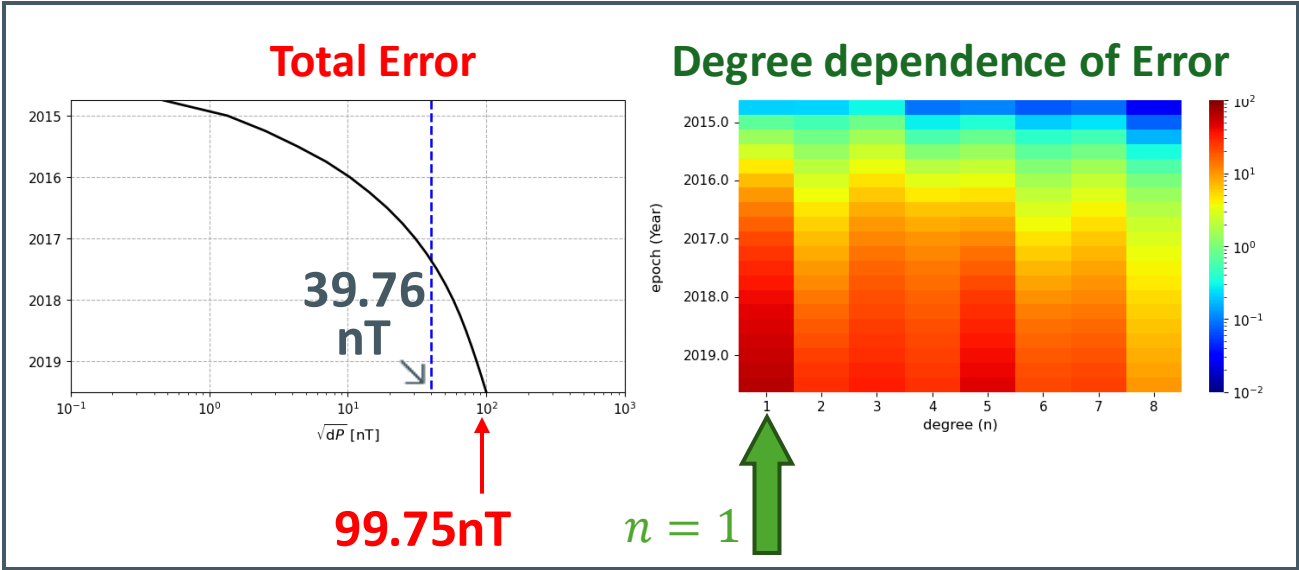
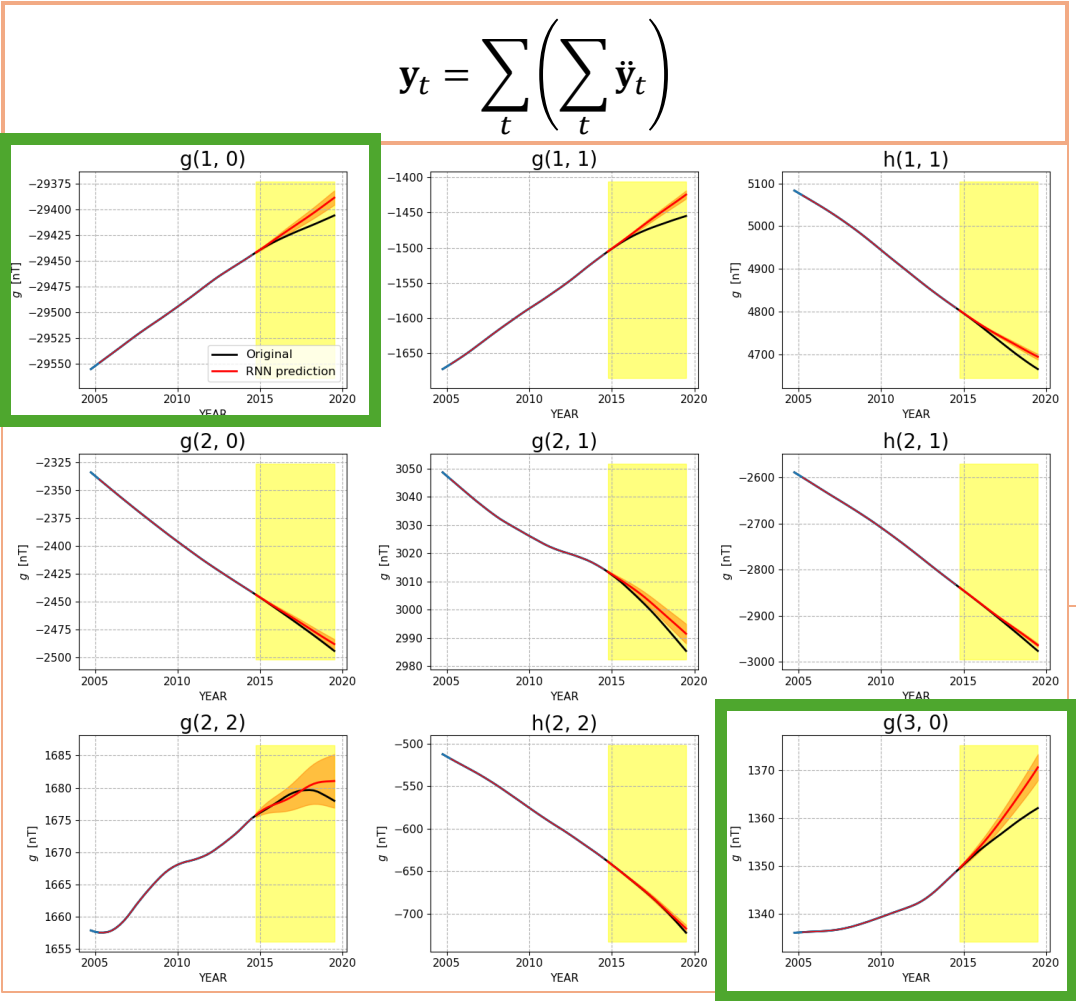
$$\dot{y}_t = \sum_t \ddot{y}_t$$



- Training Data (MCM model)
- Validation Data (MCM model)
- RNN Prediction



HINDCAST | Original time series (y)

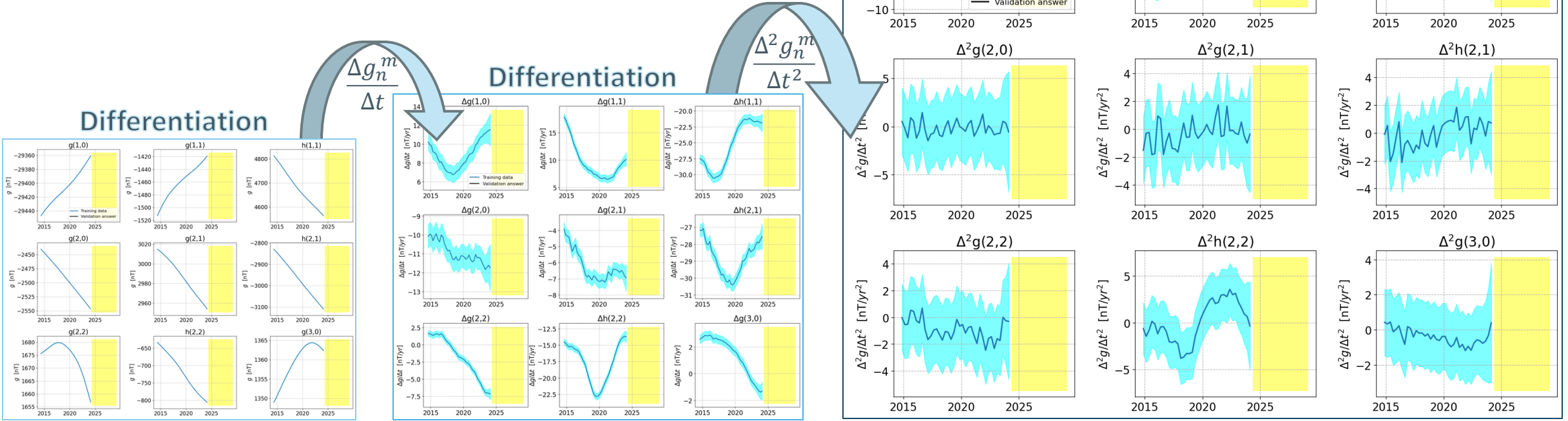


HINDCAST | Summary Table

Model	BPTT-RNN	EKF-RNN	Minami et al. (2020) (4D-EnVar)
Average $\sqrt{dP}$	33.43 – 53.69 nT	<u>37.32 – 46.62 nT</u>	---
5-year $\sqrt{dP}$	85.83 – 137.74 nT	<u>99.75 – 117.93 nT</u>	100.9 – 228.5 nT
Remarks	$\sqrt{dP}$ depends on the # of iteration of BPTT.	$\sqrt{dP}$ depends on the initial $\mathbf{w}_0$.	MHD dynamo (x1000) $\sqrt{dP}$ depends on the assim. window.

FORECAST | Forecast w/ MCM-2023

Input Data : MCM Model (Ropp et al., 2023)
Training Period : 2014.37 ~ 2024.13 (40 points)
Test Period : 2024.37 ~ 2029.13 (20 points)



FORECAST | 2nd order time derivative ($\ddot{y}$)



Common Feature:

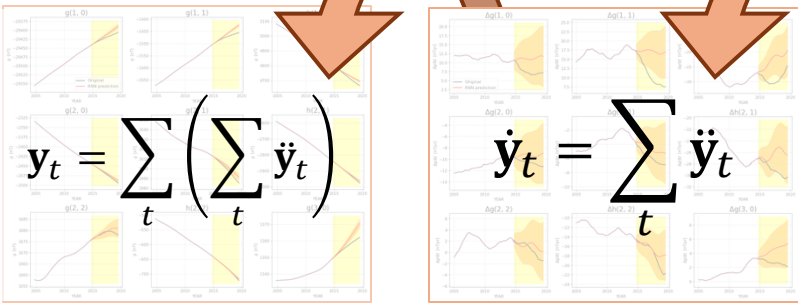
Predicted dynamics = **sin-like** wave



Strong Uncertainty

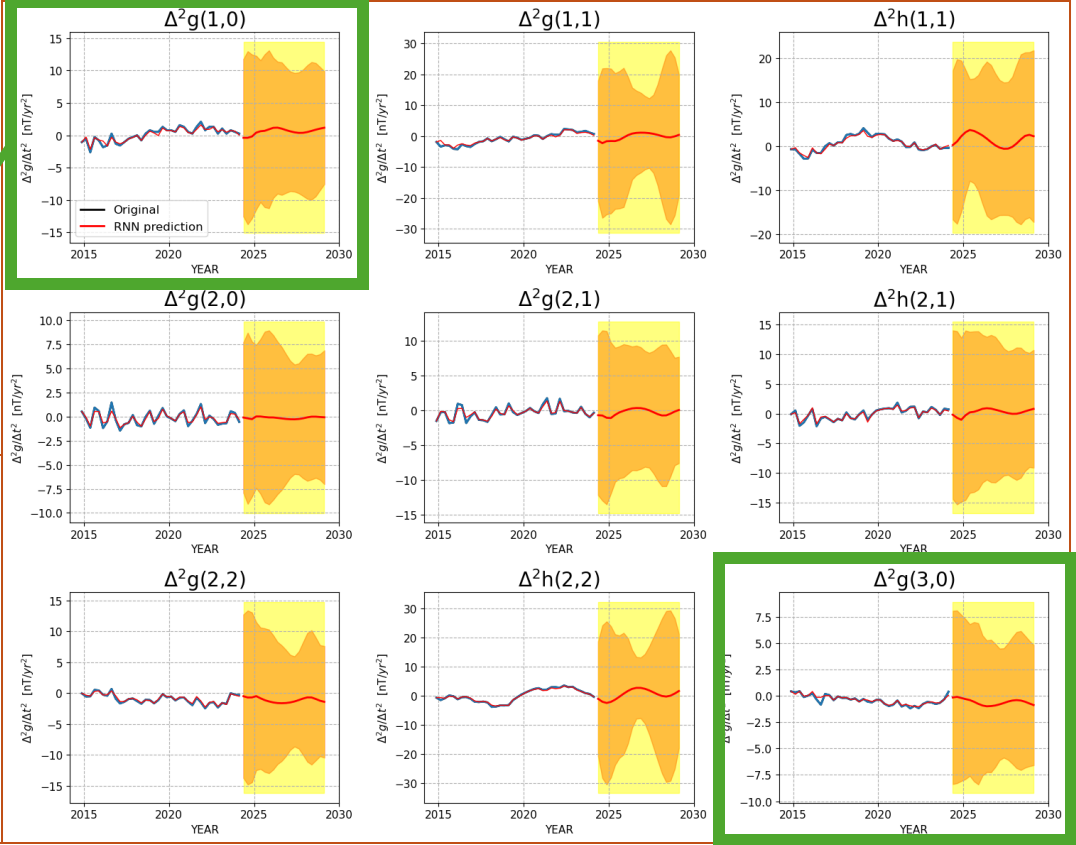
cf: $\sigma \sim 5$ [nT] in 2015-2020

Summation



- Training Data (MCM model)
- Validation Data (MCM model)
- RNN Prediction

$$\ddot{y} = \frac{\Delta^2 g_n^m}{\Delta t^2}$$



FORECAST | 1st order time derivative ($\dot{y}$)

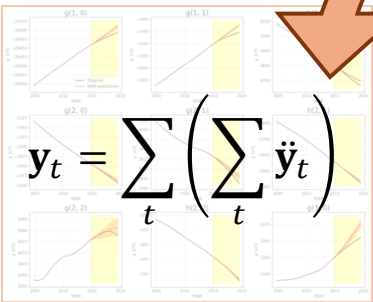


Common Feature:
~~Extends the last trend~~

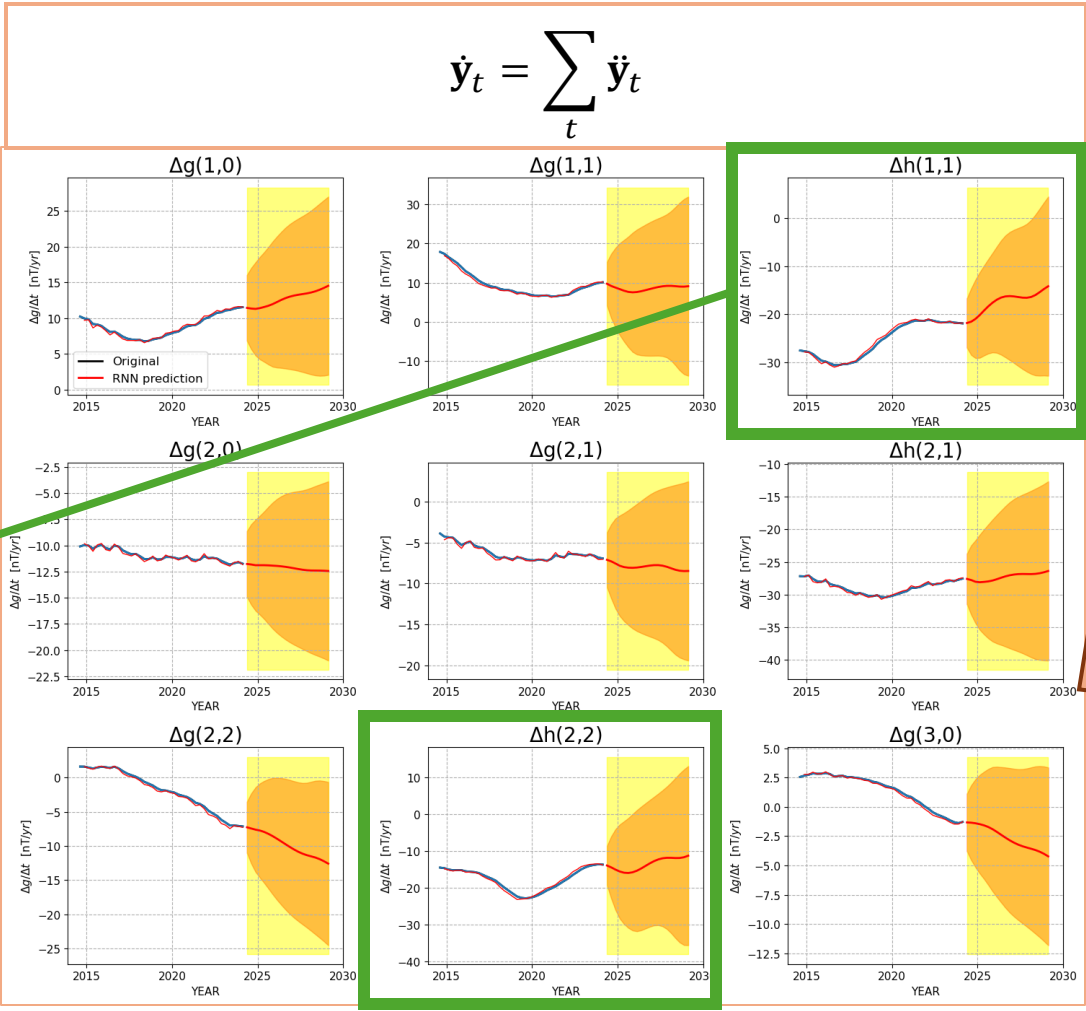


Reproduction of the former
dynamics?

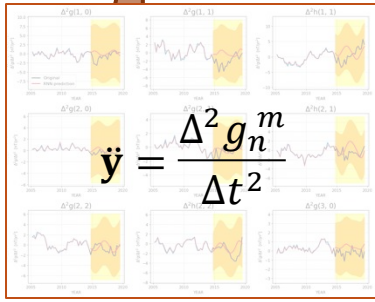
Summation



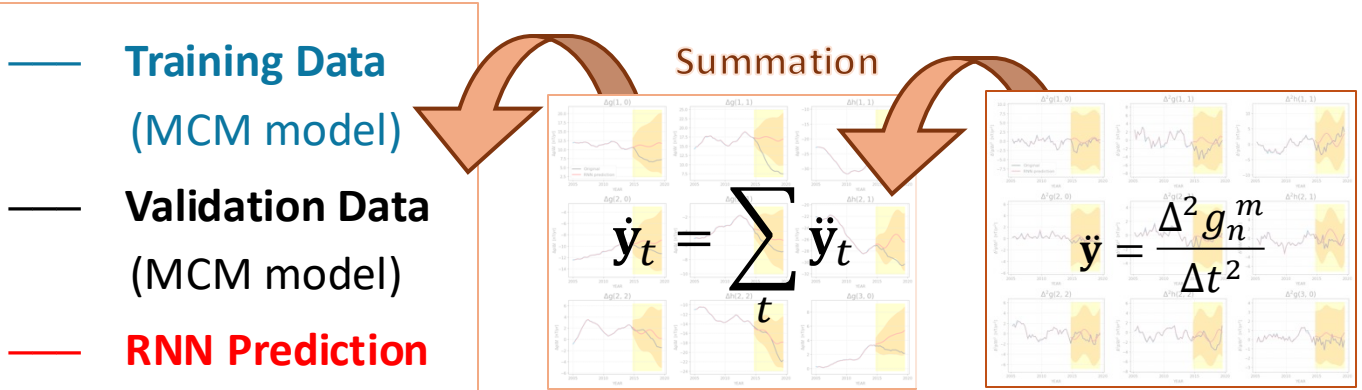
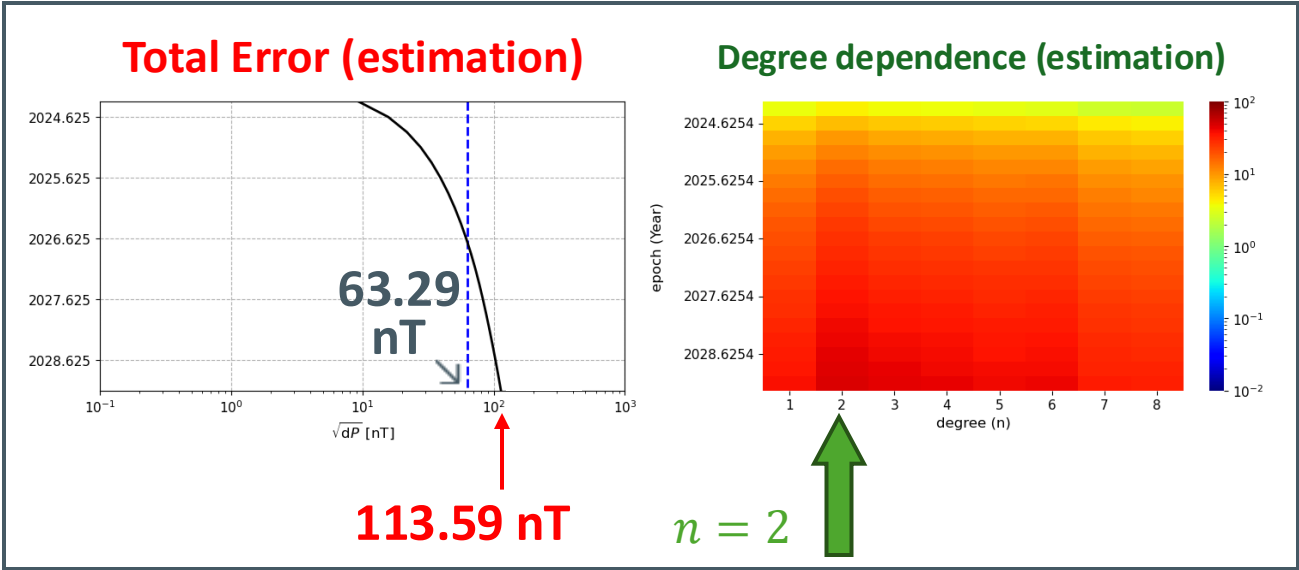
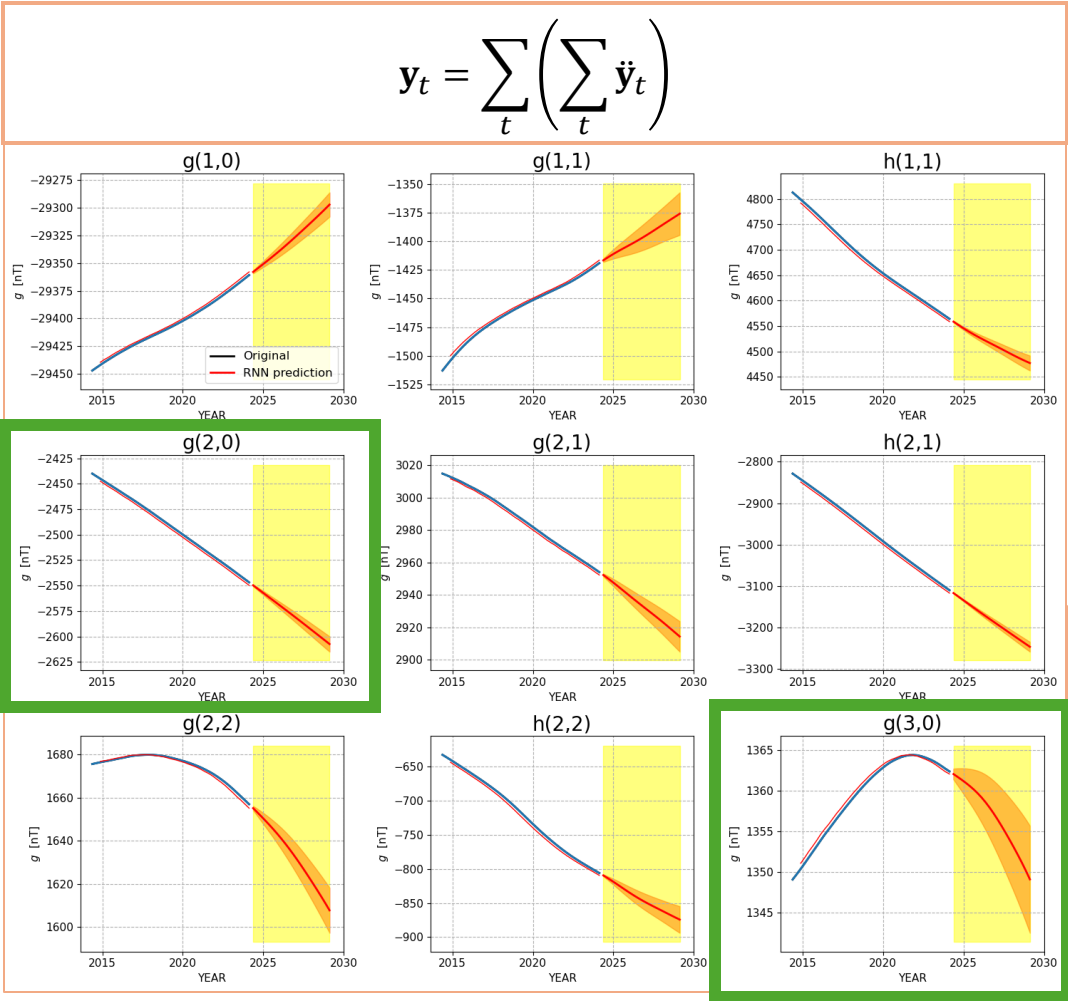
$$\dot{y}_t = \sum_t \ddot{y}_t$$



- Training Data (MCM model)
- Validation Data (MCM model)
- RNN Prediction



FORECAST | Original time series (y)



RESEARCH PLAN | Schedule & Road Map

Short-term forecast of the geomagnetic secular variation using RNNs trained by EKF

Results so far

- EKF-RNN achieves slightly better accuracy than the Geodynamo-based prediction.

ML models can predict geomagnetic variations without geodynamo models.

EKF-RNN uses error covariance matrices of training data, which makes it possible to
① calculate Forecast error - **P**, and ② Prevents **Overfitting**

Schedule & Road Map

- ✅ EKF-RNN x [Geomag **B**] → SV ± Error Hindcast
- ✅ EKF-RNN x [Geomag **B**] → SV ± Error Forecast
- 🔥 Submit pred-model to the IGRF-14 (Deadline=Oct.1)
- EKF-RNN x [Geomag **B** + Core Surface Flow **u**]

References

- Sanchez, S., Wicht, J. & Bärenzung, J. Predictions of the geomagnetic secular variation based on the ensemble sequential assimilation of geomagnetic field models by dynamo simulations. *Earth Planets Space* 72, 157 (2020). <https://doi.org/10.1186/s40623-020-01279-y>.
- Minami, T., Nakano, S., Lesur, V. et al. A candidate secular variation model for IGRF-13 based on MHD dynamo simulation and 4DEnVar data assimilation. *Earth Planets Space* 72, 136 (2020). <https://doi.org/10.1186/s40623-020-01253-8>.
- Tangborn, A., Kuang, W., Sabaka, T.J. et al. Geomagnetic secular variation forecast using the NASA GEMS ensemble Kalman filter: A candidate SV model for IGRF-13. *Earth Planets Space* 73, 47 (2021). <https://doi.org/10.1186/s40623-020-01324-w>.
- G Ropp, V Lesur, Mid-latitude and equatorial core surface flow variations derived from observatory and satellite magnetic data, *Geophysical Journal International*, Volume 234, Issue 2, August 2023, Pages 1191–1204, <https://doi.org/10.1093/gji/ggad113>.
- Ropp, G., Lesur, V., Baerenzung, J. et al. Sequential modelling of the Earth's core magnetic field. *Earth Planets Space* 72, 153 (2020). <https://doi.org/10.1186/s40623-020-01230-1>.
- Puskorius GV, Feldkamp LA. Neurocontrol of nonlinear dynamical systems with Kalman filter trained recurrent networks. *IEEE Trans Neural Netw.* 1994;5(2):279-97. <https://doi.org/10.1109/72.279191>.
- Elman, Jeffrey L. (1990). "Finding Structure in Time". *Cognitive Science*. 14 (2): 179–211. https://doi.org/10.1207/s15516709cog1402_1.

METHOD | Differentiating training data $\ddot{\mathbf{y}}$ & $\ddot{\mathbf{R}}$

MCM-model: $\mathbf{y}^O = [g_0^1, g_1^1, h_1^1, \dots, h_{13}^{13}]$, $\mathbf{R} = \text{Obs-Error Matrix}$, $(\Delta t = 0.25 \text{ year})$

--> 1st deriv: $\dot{\mathbf{y}}^O(t_i) = \frac{\mathbf{y}^O(t_i) - \mathbf{y}^O(t_{i-1})}{\Delta t}$, $\dot{\mathbf{R}}(t_i) = \frac{1}{\Delta t^2} (\mathbf{R}(t_i) + \mathbf{R}(t_{i-1}))$

--> 2nd deriv: $\ddot{\mathbf{y}}^O(t_i) = \frac{\mathbf{y}^O(t_i) - 2\mathbf{y}^O(t_{i-1}) + \mathbf{y}^O(t_{i-2}))}{\Delta t^2}$, $\ddot{\mathbf{R}}(t_i) = \frac{1}{\Delta t^4} (\mathbf{R}(t_i) + 4\mathbf{R}(t_{i-1}) + \mathbf{R}(t_{i-2}))$

P-inflation (to prevent filter divergence): RTPS*1 $\mathbf{P}^{\text{inf}}(t) = \alpha \mathbf{P}^f(t) + (1 - \alpha) \mathbf{P}^a(t)$ ($\alpha = 0.5$)

Forecast Loop

2nd deriv: $\mathbf{z}(t_{i+1}) = \text{RNN}(\mathbf{z}(t_i)) = \ddot{\mathbf{y}}^F$, $\mathbf{P}_z(\mathbf{t}_{i+1}) = \mathbf{G}(t_i) \mathbf{P}_z(t_i) \mathbf{G}(t_i)^T + \mathbf{H}(t_i) \mathbf{P}^a(T) \mathbf{H}(t_i)^T = \ddot{\mathbf{R}}^F$

--> 1st deriv: $\dot{\mathbf{y}}^F(t_{i+1}) = \dot{\mathbf{y}}^F(t_i) + \ddot{\mathbf{y}}^F(t_i) \cdot \Delta t$, $\dot{\mathbf{R}}^F(t_{i+1}) = \dot{\mathbf{R}}^F(t_i) + \ddot{\mathbf{R}}^F(t_i) \cdot \Delta t^2$

--> Target: $\mathbf{y}^F(t_{i+1}) = \mathbf{y}^F(t_i) + \dot{\mathbf{y}}^F(t_i) \cdot \Delta t$, $\mathbf{R}^F(t_{i+1}) = \mathbf{R}^F(t_i) + \dot{\mathbf{R}}^F(t_i) \cdot \Delta t^2 + \ddot{\mathbf{R}}^F(t_i) \cdot \Delta t^4$

$$\left(\mathbf{L}_z(t, \mathbf{x}) = \frac{\partial \mathbf{z}}{\partial \mathbf{x}} \Big|_{\mathbf{w}=\mathbf{w}(t)} , \quad \mathbf{L}_w(t, \mathbf{w}) = \frac{\partial \mathbf{z}}{\partial \mathbf{w}} \Big|_{\mathbf{x}=\mathbf{x}(t)} , \quad \mathbf{G}(t) = \mathbf{L}_z(t, \mathbf{x}(t)) , \quad \mathbf{H}(t) = \mathbf{L}_w(t, \mathbf{w}(t)) \right)$$

*1: Relaxation-to-Prior-Spread method (Whitaker and Hamill, 2012)

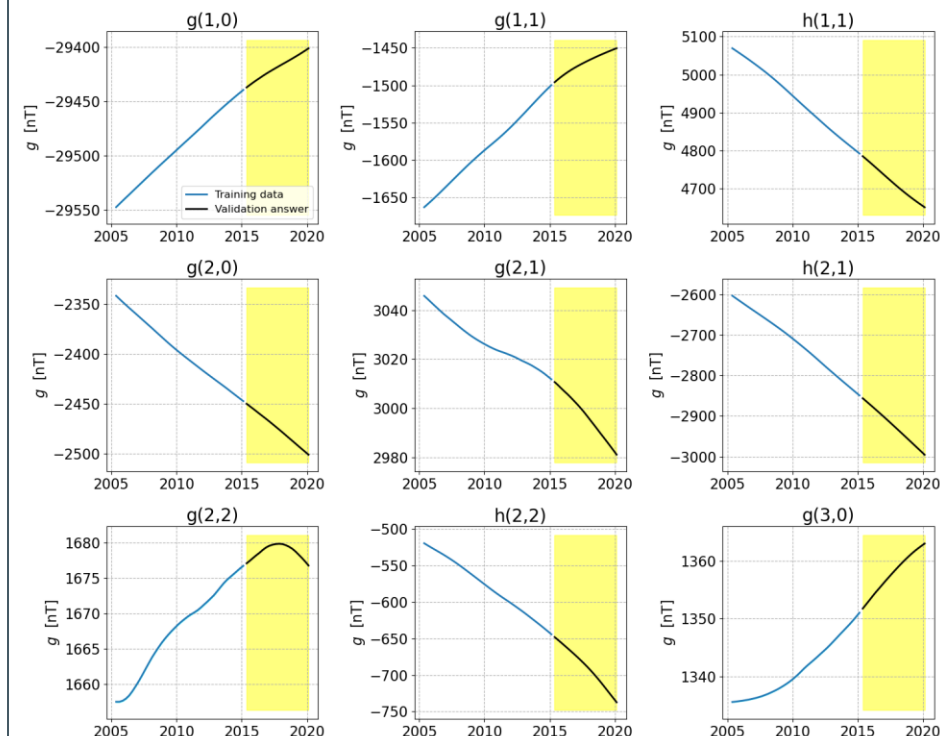
DETAIL | Differentiating training data

$$\mathbf{y}^O(t_i) = [g_1^0(t_i), g_1^1(t_i), h_1^1(t_i), \dots, g_{13}^{13}(t_i), h_{13}^{13}(t_i)].$$

$$\boldsymbol{\varepsilon}^O(t_i) = \mathbf{y}^O(t_i) - \mathbf{x}^{\text{true}}(t_i).$$

$$\mathbf{R}(t_i) = \langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_i)^T \rangle$$

$$= \begin{bmatrix} \delta g_1^0(t_i)^2 & \dots & \langle \text{Cross Term} \rangle \\ \vdots & \ddots & \vdots \\ \langle \text{Cross Term} \rangle & \dots & \delta h_{13}^{13}(t_i)^2 \end{bmatrix}.$$



DETAIL | Differentiating training data

$$\dot{\mathbf{y}}^O(t_i) = \frac{\mathbf{y}^O(t_i) - \mathbf{y}^O(t_{i-1})}{\Delta t}.$$

(Backward diff.
 $\Delta t = 0.25$ year)

$$\dot{\mathbf{x}}^{\text{true}}(t_i) = \frac{\mathbf{x}^{\text{true}}(t_i) - \mathbf{x}^{\text{true}}(t_{i-1})}{\Delta t}.$$

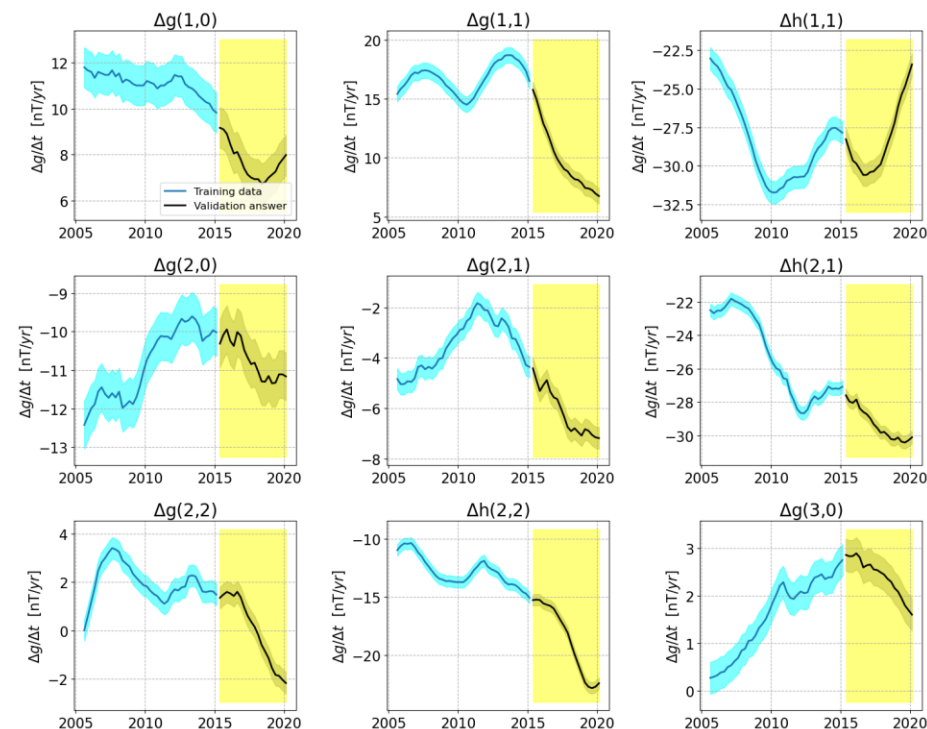
$$\dot{\boldsymbol{\varepsilon}}^O(t_i) = \dot{\mathbf{y}}^O(t_i) - \dot{\mathbf{x}}^{\text{true}}(t_i) = \frac{1}{\Delta t} \left(\boldsymbol{\varepsilon}^O(t_i) - \boldsymbol{\varepsilon}^O(t_{i-1}) \right).$$

$$\dot{\mathbf{R}}(t_i) = \langle \dot{\boldsymbol{\varepsilon}}^O(t_i) \cdot \dot{\boldsymbol{\varepsilon}}^O(t_i)^T \rangle$$

$$= \frac{1}{\Delta t^2} \left(\begin{aligned} &\langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_i)^T \rangle + \langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_{i-1})^T \rangle \\ &+ \langle \boldsymbol{\varepsilon}^O(t_{i-1}) \cdot \boldsymbol{\varepsilon}^O(t_i)^T \rangle + \langle \boldsymbol{\varepsilon}^O(t_{i-1}) \cdot \boldsymbol{\varepsilon}^O(t_{i-1})^T \rangle \end{aligned} \right).$$

Assume $\langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_{i-1})^T \rangle = \langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_{i-1})^T \rangle = 0$,

$$\dot{\mathbf{R}}(t_i) = \mathbf{R}^O(t_i) + \mathbf{R}^O(t_{i-1})$$



DETAIL | Differentiating training data

$$\ddot{\mathbf{y}}^O(t_i) = \frac{\dot{\mathbf{y}}^O(t_i) - \dot{\mathbf{y}}^O(t_{i-1})}{\Delta t} = \frac{\mathbf{y}^O(t_i) - 2\mathbf{y}^O(t_{i-1}) + \mathbf{y}^O(t_{i-2}))}{\Delta t^2}.$$

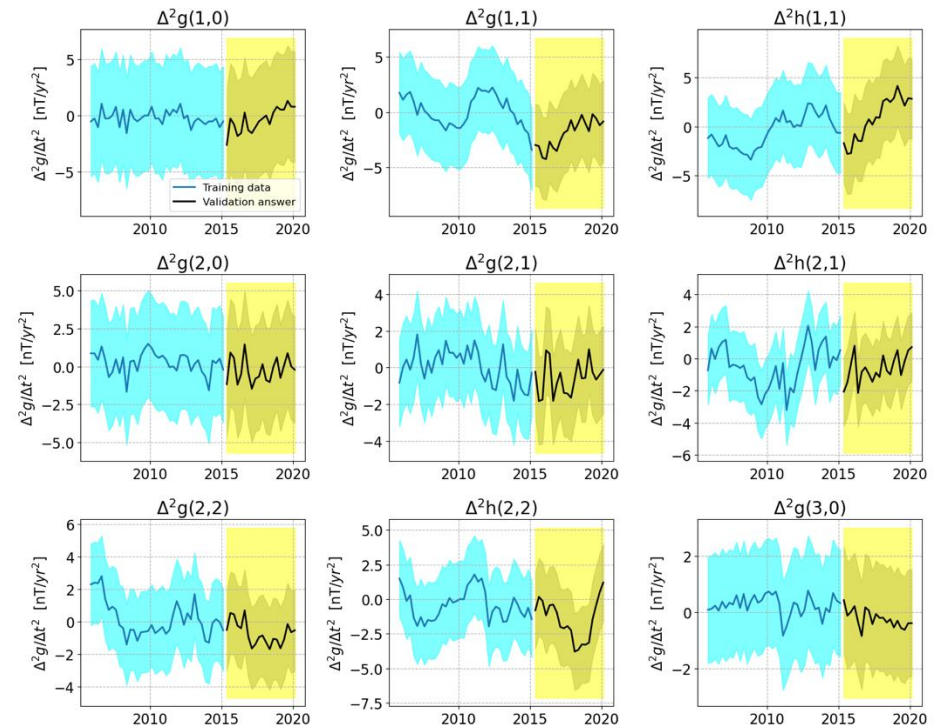
$$\ddot{\boldsymbol{\varepsilon}}^O(t_i) = \ddot{\mathbf{y}}^O(t_i) - \ddot{\mathbf{x}}^{\text{true}}(t_i) = \frac{1}{\Delta t^2} \left(\boldsymbol{\varepsilon}^O(t_i) - 2\boldsymbol{\varepsilon}^O(t_{i-1}) + \boldsymbol{\varepsilon}^O(t_{i-2}) \right).$$

$$\ddot{\mathbf{R}}(t_i) = \langle \ddot{\boldsymbol{\varepsilon}}^O(t_i) \cdot \ddot{\boldsymbol{\varepsilon}}^O(t_i)^T \rangle$$

$$= \frac{1}{\Delta t^4} \begin{pmatrix} \langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_i)^T \rangle - 2 \langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_{i-1})^T \rangle + \langle \boldsymbol{\varepsilon}^O(t_i) \cdot \boldsymbol{\varepsilon}^O(t_{i-2})^T \rangle \\ -2 \langle \boldsymbol{\varepsilon}^O(t_{i-1}) \cdot \boldsymbol{\varepsilon}^O(t_i)^T \rangle + 4 \langle \boldsymbol{\varepsilon}^O(t_{i-1}) \cdot \boldsymbol{\varepsilon}^O(t_{i-1})^T \rangle - 2 \langle \boldsymbol{\varepsilon}^O(t_{i-1}) \cdot \boldsymbol{\varepsilon}^O(t_{i-2})^T \rangle \\ + \langle \boldsymbol{\varepsilon}^O(t_{i-2}) \cdot \boldsymbol{\varepsilon}^O(t_i)^T \rangle - 2 \langle \boldsymbol{\varepsilon}^O(t_{i-2}) \cdot \boldsymbol{\varepsilon}^O(t_{i-1})^T \rangle + \langle \boldsymbol{\varepsilon}^O(t_{i-2}) \cdot \boldsymbol{\varepsilon}^O(t_{i-2})^T \rangle \end{pmatrix}.$$

$$\text{Assume } \langle \boldsymbol{\varepsilon}^O(t_j) \cdot \boldsymbol{\varepsilon}^O(t_k)^T \rangle = \begin{cases} \langle \boldsymbol{\varepsilon}^O(t_j)^2 \rangle & (\text{for } j = k) \\ 0 & (\text{for } j \neq k) \end{cases},$$

$$\dot{\mathbf{R}}(t_i) = \mathbf{R}^O(t_i) + 4\mathbf{R}^O(t_{i-1}) + \mathbf{R}^O(t_{i-2})$$



DETAIL | Integrating forecast data

$$\ddot{\mathbf{y}}^F(t_i) = \text{RNN}(\ddot{\mathbf{y}}^F(t_{i-1})), \quad \mathbf{P}_z(t_i) = \mathbf{G}(t_{i-1})\mathbf{P}_z(t_{i-1})\mathbf{G}(t_{i-1})^T + \mathbf{H}(t_{i-1})\mathbf{P}^a(t_{\text{Rel}})\mathbf{H}(t_{i-1})^T = \ddot{\mathbf{R}}^F(t_i).$$

$$\begin{aligned} \dot{\mathbf{y}}^F(t_i) &= \dot{\mathbf{y}}^F(t_{i-1}) + \ddot{\mathbf{y}}^F(t_i) \cdot \Delta t \\ \dot{\mathbf{x}}^{\text{true}}(t_i) &= \dot{\mathbf{x}}^{\text{true}}(t_{i-1}) + \ddot{\mathbf{x}}^{\text{true}}(t_i) \cdot \Delta t \end{aligned}$$

$$\dot{\boldsymbol{\varepsilon}}^F(t_i) = \dot{\boldsymbol{\varepsilon}}^F(t_{i-1}) + \ddot{\boldsymbol{\varepsilon}}^F(t_i) \cdot \Delta t$$

$$\text{Assume } \langle \dot{\boldsymbol{\varepsilon}}^F(t_i) \cdot \ddot{\boldsymbol{\varepsilon}}^F(t_i)^T \rangle = \langle \ddot{\boldsymbol{\varepsilon}}^F(t_i) \cdot \dot{\boldsymbol{\varepsilon}}^F(t_i)^T \rangle = 0$$

$$\dot{\mathbf{R}}^F(t_i) = \dot{\mathbf{R}}^F(t_{i-1}) + \ddot{\mathbf{R}}^F(t_{i-1}) \cdot \Delta t^2$$

$$\begin{aligned} \mathbf{y}^F(t_i) &= \mathbf{y}^F(t_{i-1}) + \dot{\mathbf{y}}^F(t_i) \cdot \Delta t \\ \mathbf{x}^{\text{true}}(t_i) &= \mathbf{x}^{\text{true}}(t_{i-1}) + \dot{\mathbf{x}}^{\text{true}}(t_i) \cdot \Delta t \end{aligned}$$

$$\begin{aligned} \boldsymbol{\varepsilon}^F(t_i) &= \boldsymbol{\varepsilon}^F(t_{i-1}) + \dot{\boldsymbol{\varepsilon}}^F(t_i) \cdot \Delta t \\ &= \boldsymbol{\varepsilon}^F(t_{i-1}) + \dot{\boldsymbol{\varepsilon}}^F(t_{i-1}) \cdot \Delta t + \ddot{\boldsymbol{\varepsilon}}^F(t_i) \cdot \Delta t^2 \end{aligned}$$

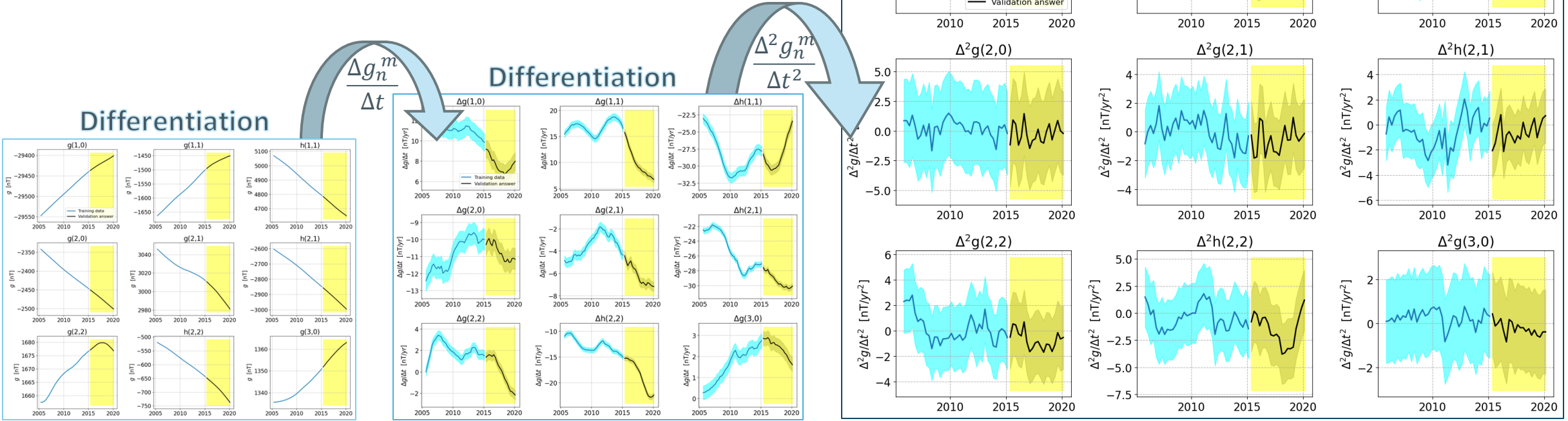
$$\begin{aligned} \text{Assume } \langle \boldsymbol{\varepsilon}^F(t_i) \cdot \dot{\boldsymbol{\varepsilon}}^F(t_i)^T \rangle &= \langle \dot{\boldsymbol{\varepsilon}}^F(t_i) \cdot \boldsymbol{\varepsilon}^F(t_i)^T \rangle \\ \langle \boldsymbol{\varepsilon}^F(t_i) \cdot \ddot{\boldsymbol{\varepsilon}}^F(t_i)^T \rangle &= \langle \ddot{\boldsymbol{\varepsilon}}^F(t_i) \cdot \boldsymbol{\varepsilon}^F(t_i)^T \rangle = 0 \end{aligned}$$

$$\mathbf{R}^F(t_i) = \mathbf{R}^F(t_{i-1}) + \dot{\mathbf{R}}^F(t_{i-1}) \cdot \Delta t^2 + \ddot{\mathbf{R}}^F(t_{i-1}) \cdot \Delta t^4$$

METHOD | Hindcast w/ MCM-2023

We assess **RNNs** under the same evaluation setup as in ~~Minami et al. (2020)~~

Input Data : MCM Model (Ropp et al., 2023)
Training Period : 2005.37 ~ 2015.13 (40 points)
Validation Period : 2015.63 ~ 2020.13 (20 points)



HINDCAST | 2nd order time derivative ($\ddot{y}$)



Common Feature:

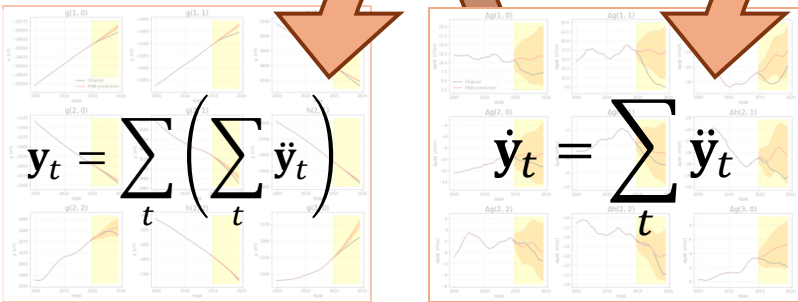
Predicted dynamics = **sin-like** wave



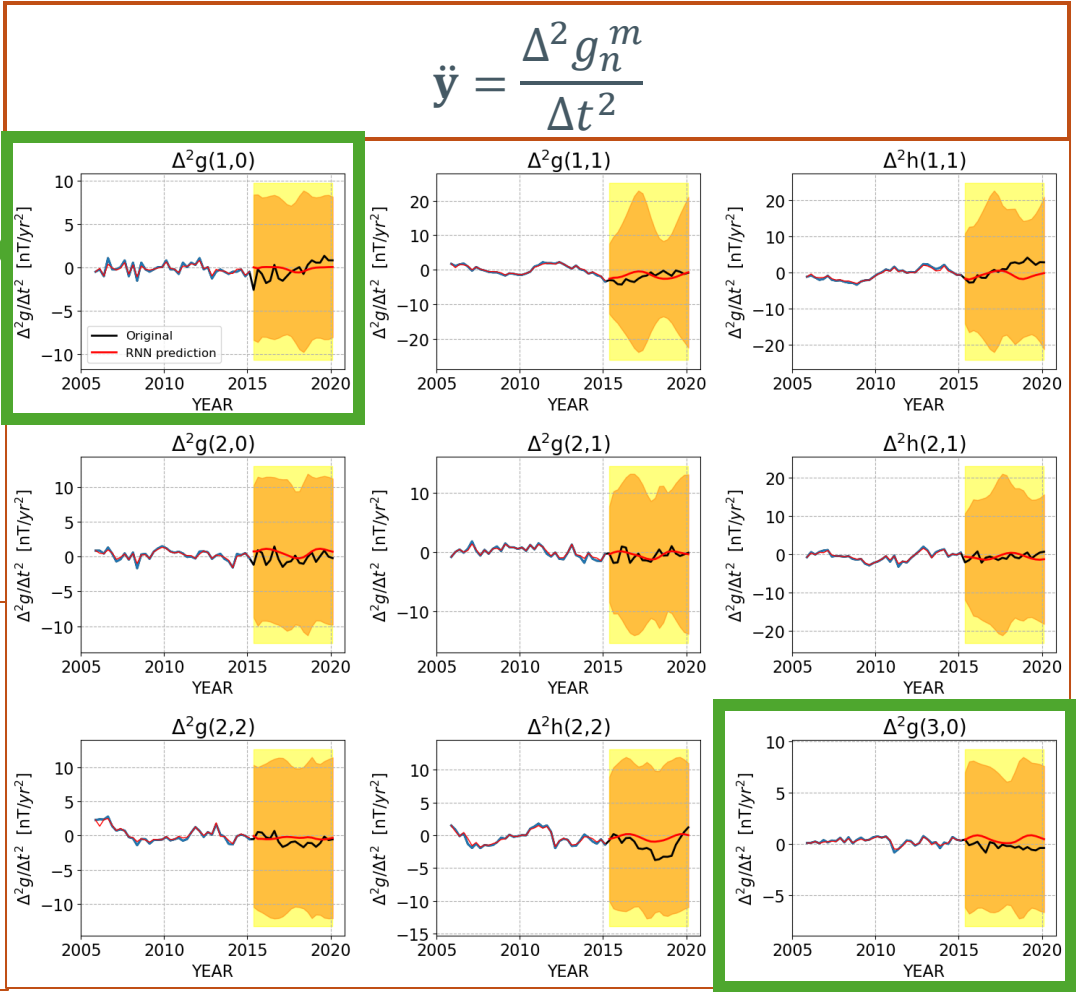
Inaccurate Component (Ex. g_1^0, g_3^0):

Incorrect positive and negative sign

Summation



- Training Data (MCM model)
- Validation Data (MCM model)
- RNN Prediction



HINDCAST | 1st order time derivative ($\dot{y}$)



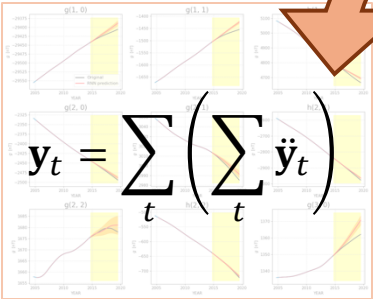
Common Feature:
~~Extends the last trend~~



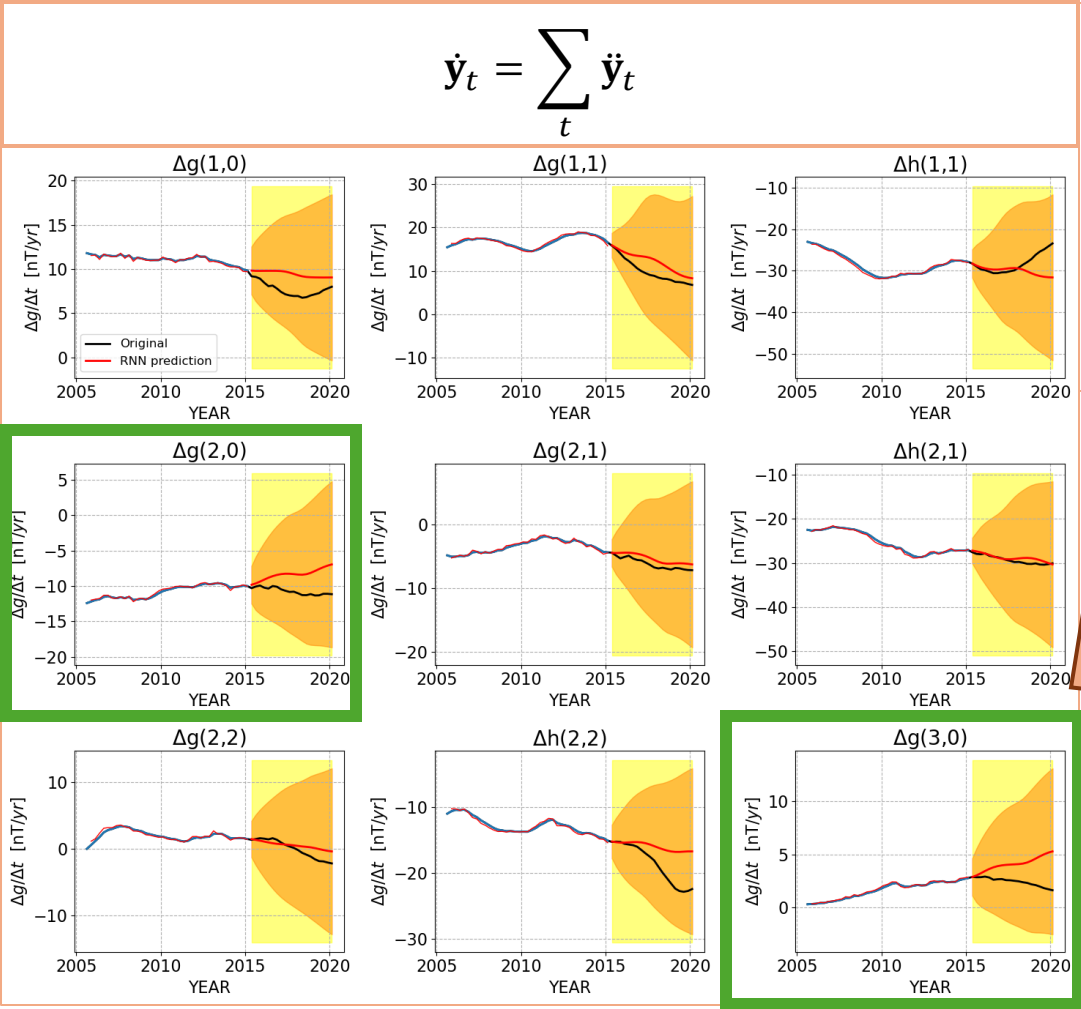
(Ex. g_2^0, g_3^0)
Inaccurate Component:

Inaccurate forecast caused by the sudden change of the trend in the validation period (2014).

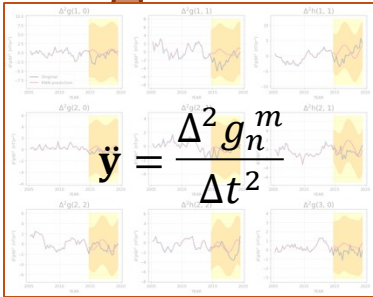
Summation



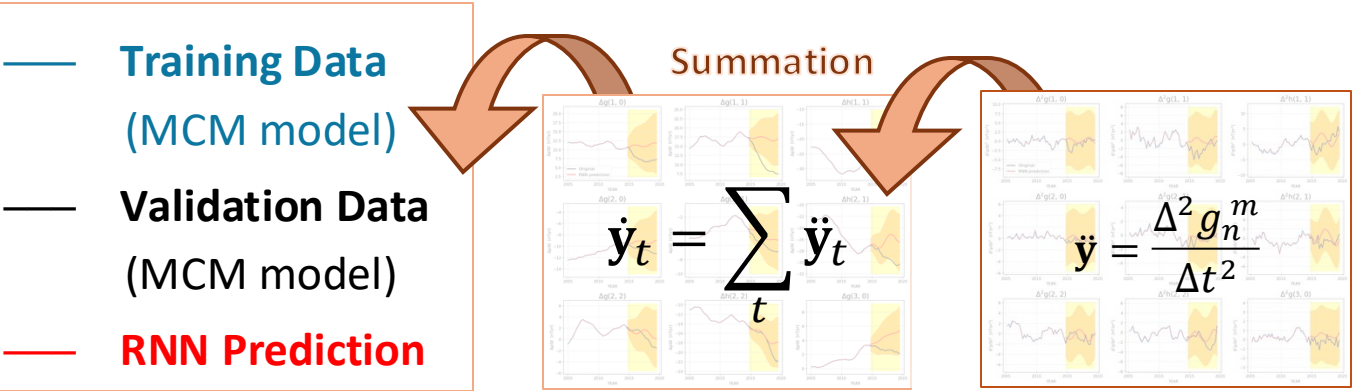
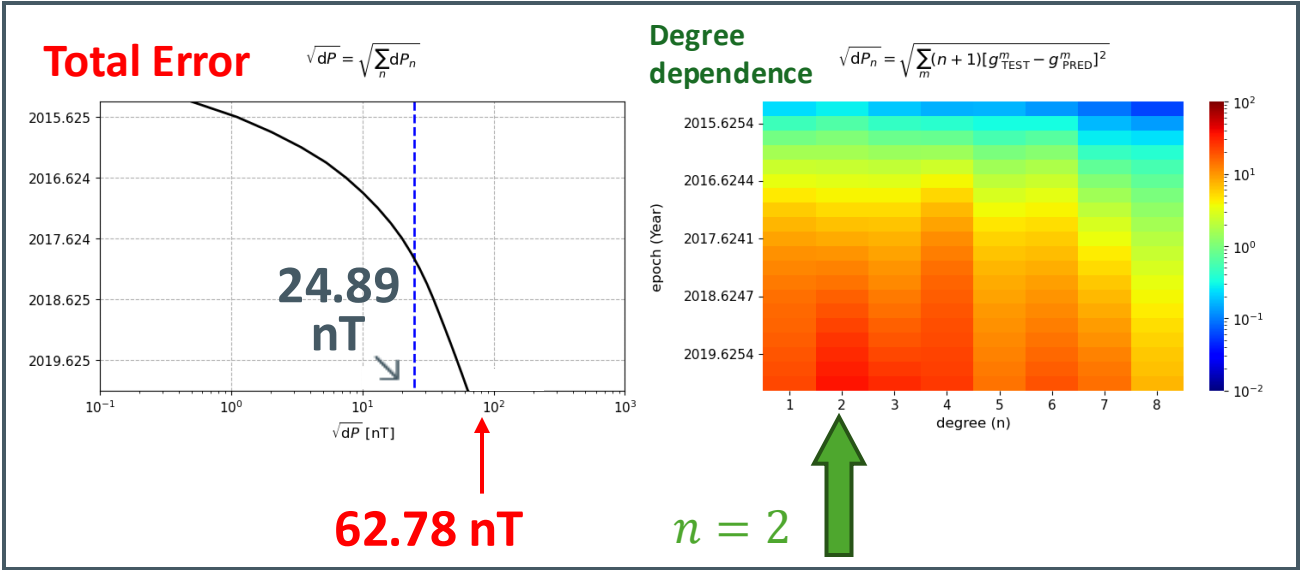
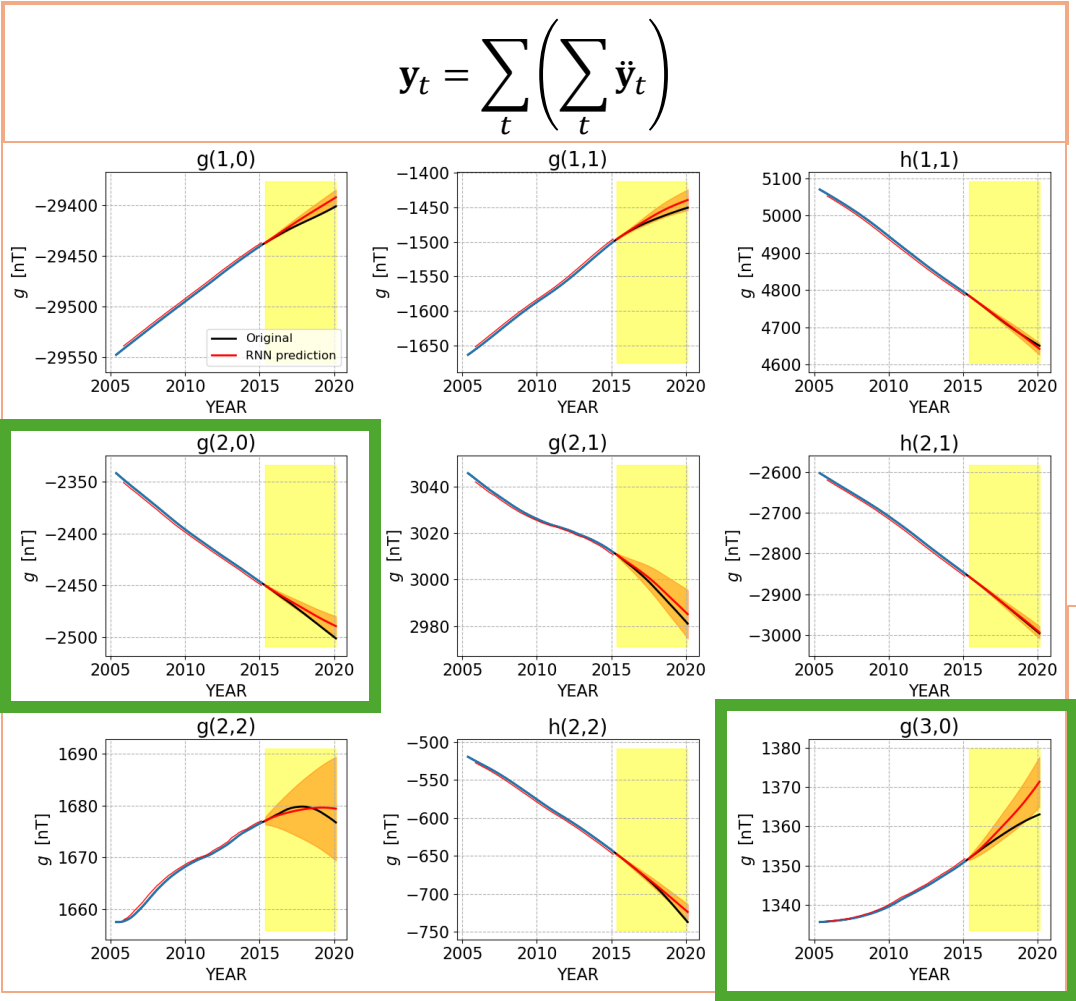
$$\dot{y}_t = \sum_t \ddot{y}_t$$



- Training Data (MCM model)
- Validation Data (MCM model)
- RNN Prediction



HINDCAST | Original time series (y)



HINDCAST | RNN hidden layer

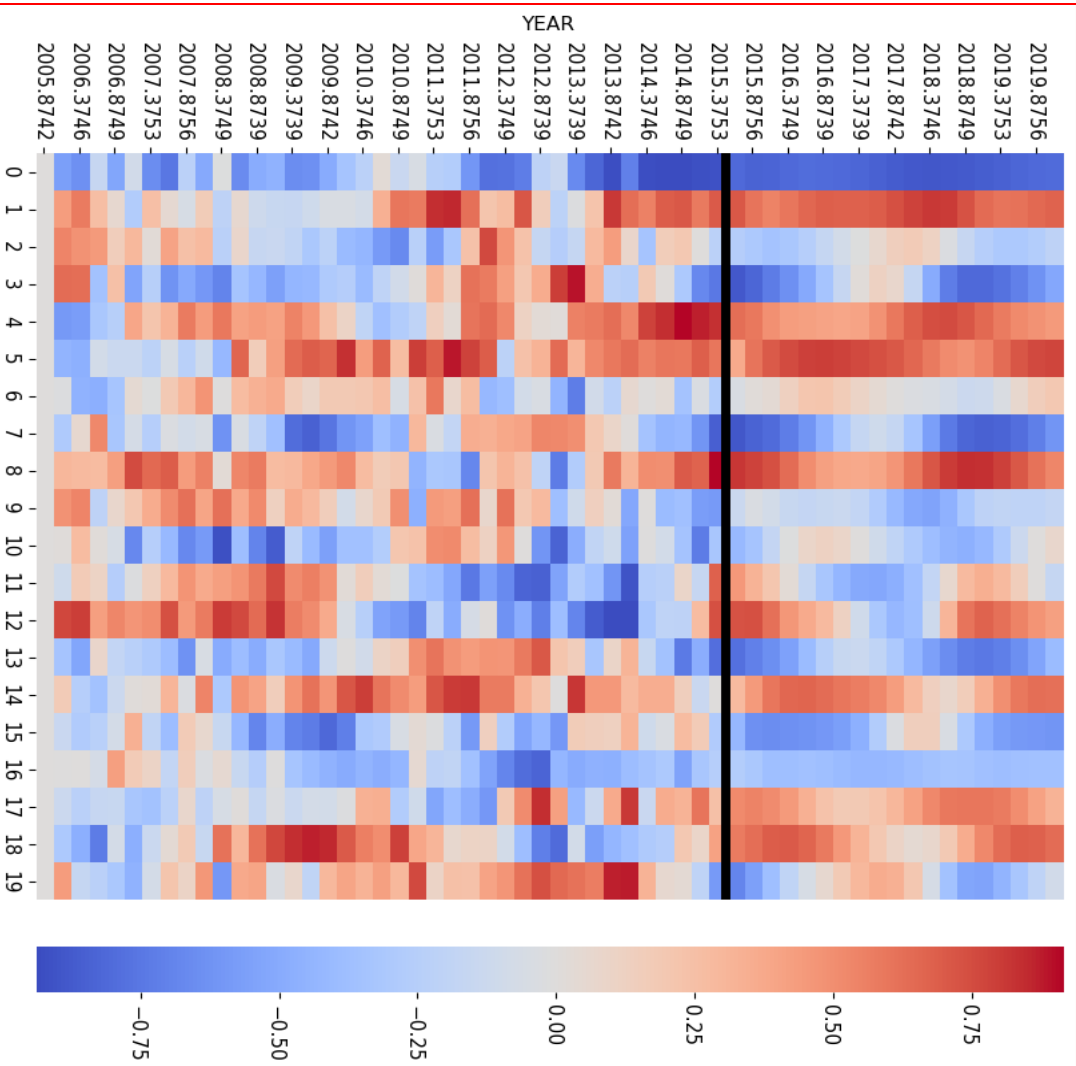
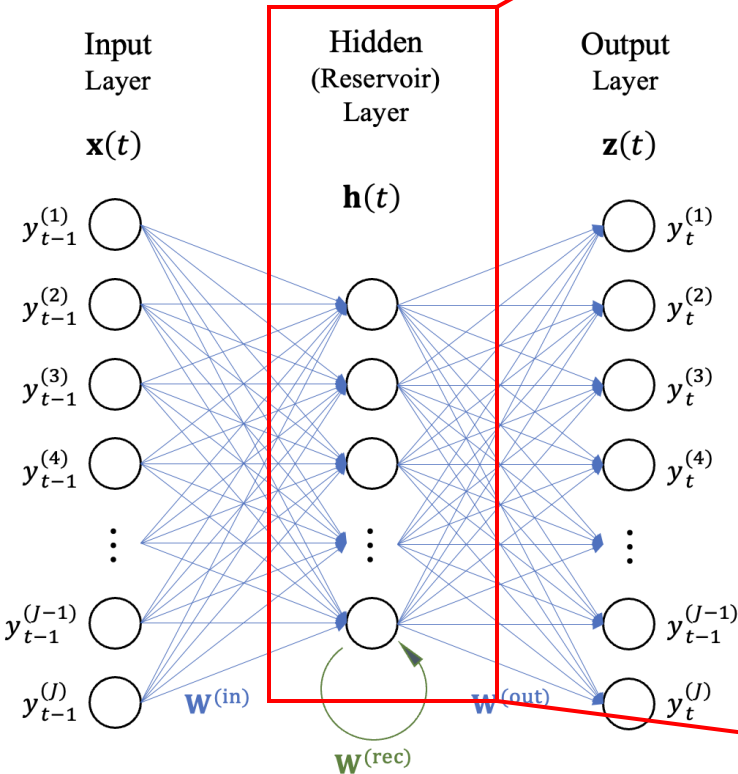
Elman Network

RNN cell

$$\mathbf{h}(t) = \tanh(\mathbf{W}^{(\text{in})} \cdot \mathbf{x}(t) + \mathbf{W}^{(\text{rec})} \cdot \mathbf{h}(t - 1) + \mathbf{b}^{(\text{rec})})$$

Affine (Dense)

$$\mathbf{z}(t) = \mathbf{W}^{(\text{out})} \cdot \mathbf{h}(t) + \mathbf{b}^{(\text{out})}$$



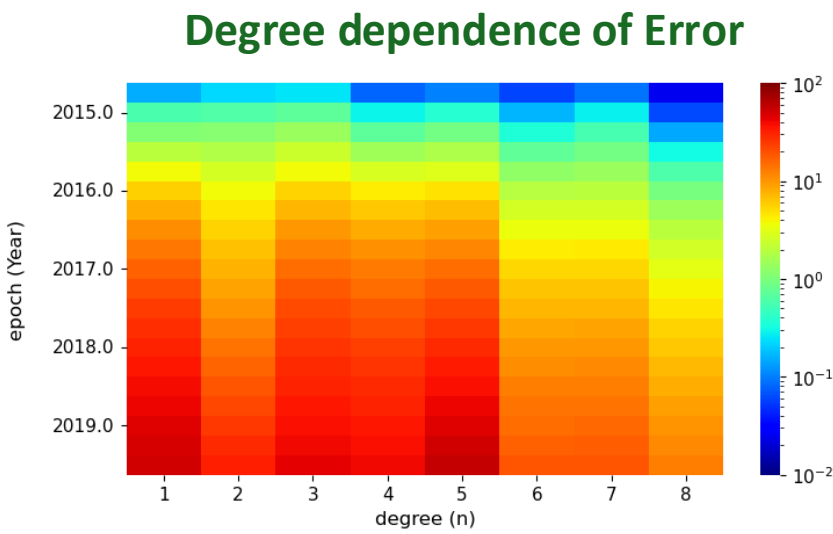
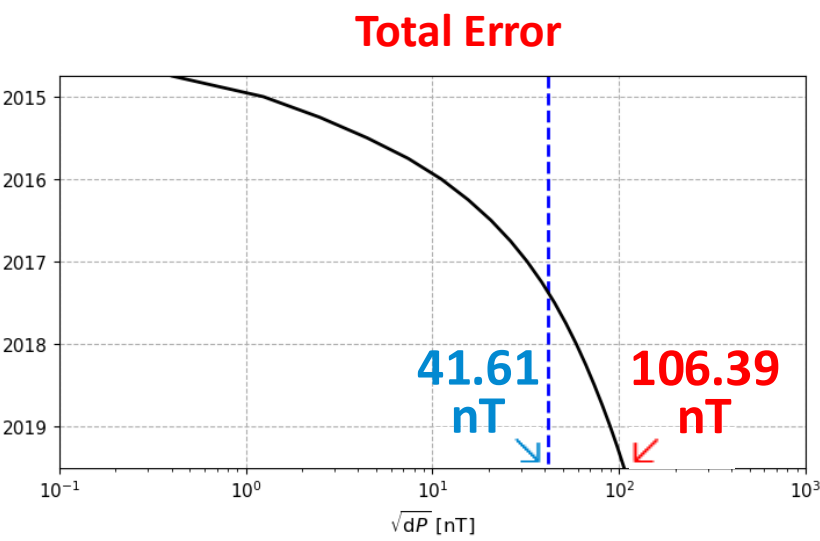
Appendix | Gauss coefs / Forecast Error Chart

$$\begin{cases} \text{rot} \vec{B} = \vec{0} \\ \text{div} \vec{B} = 0 \end{cases} \iff \begin{cases} \vec{B} = -\text{grad } U \\ \text{div grad } U = -\nabla^2 U = 0 \end{cases}$$

$$U(r, \theta, \phi) = R_E \sum_{n=1}^{\infty} \sum_{m=0}^n \left\{ (g_n^m \cos m\phi + h_n^m \sin m\phi) \left(\frac{R_E}{r} \right)^{n+1} \right\} P_n^m(\cos \theta) + R_E \sum_{n=1}^{\infty} \sum_{m=0}^n \left\{ (q_n^m \cos m\phi + s_n^m \sin m\phi) \left(\frac{r}{R_E} \right)^n \right\} P_n^m(\cos \theta)$$

Degree dependence of Error : $\sqrt{dP_n} = \sqrt{\sum_m (n+1) [g_{\text{TEST}}^m - g_{\text{PRED}}^m]^2}$

Total Error : $\sqrt{dP} = \sqrt{\sum_n dP_n}$



Mauersberger-Lowes spectrum
(Langel and Estes, 1982)

