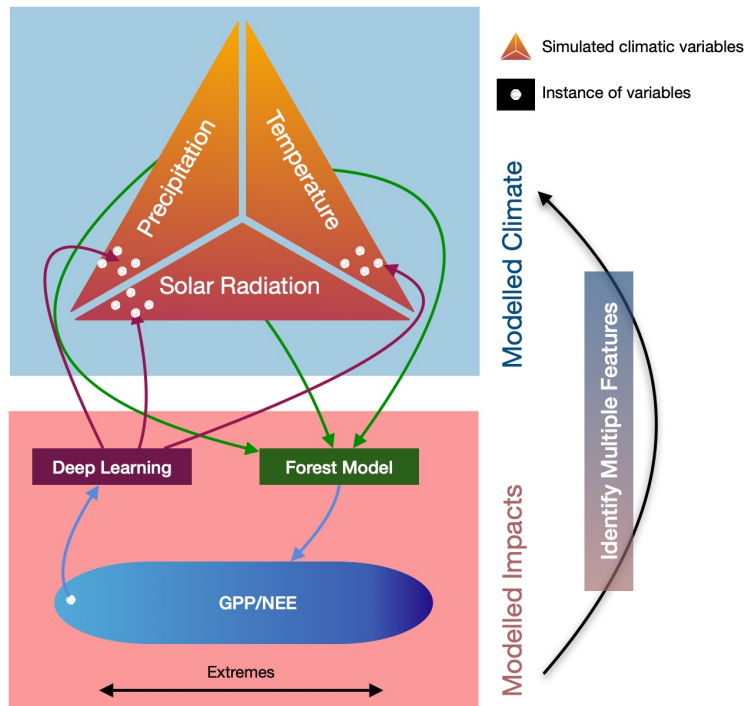




Identifying drivers of extreme reduction in carbon uptake of forests using interpretable machine learning

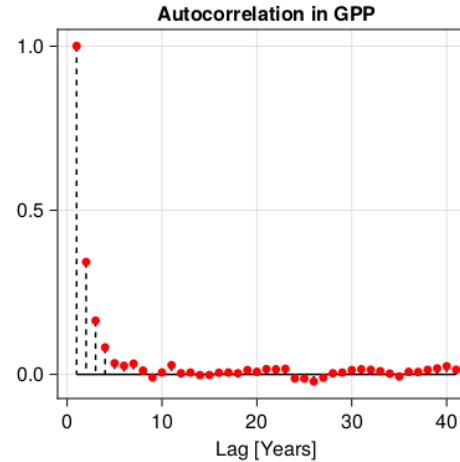
Mohit Anand¹, Gustau Camps-Valls², Jakob Zscheischler¹

EGU 2022 | 23.05.2022

1. Helmholtz Centre for Environmental Research – UFZ, Leipzig, Germany
2. Image Processing Laboratory (IPL), Universitat de València, València, Spain

Location and data

- Hainich beech forest, Germany
- Weather generator (AWE-GEN)
- 200,000 years of climate simulation
 - Precipitation
 - Temperature
 - Solar radiation
- Temporal resolution - daily
- Forest model (FORMIND)
- GPP as yearly output



- Deep Learning input-output (X,Y) pair
 - Input = $4 \times 365 \times 3$
 - Output = 1

Neural Network Architecture

- Self Explanatory Neural Network (SENN)

$$f(x) = \sum_i^k \theta(x)_i h(x)_i$$

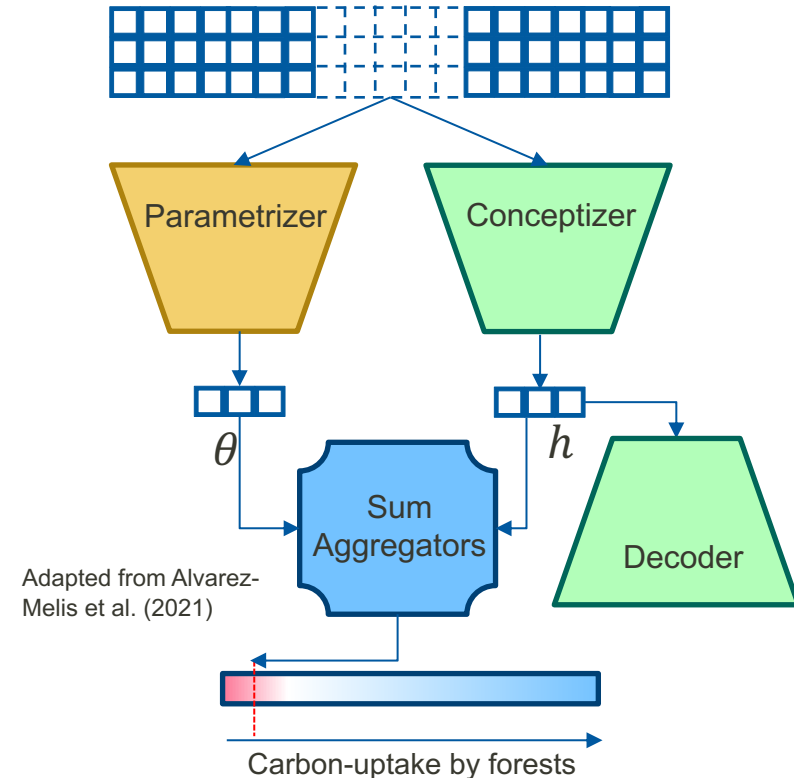
- For small changes in concepts, parameters don't change much.

- Loss Function

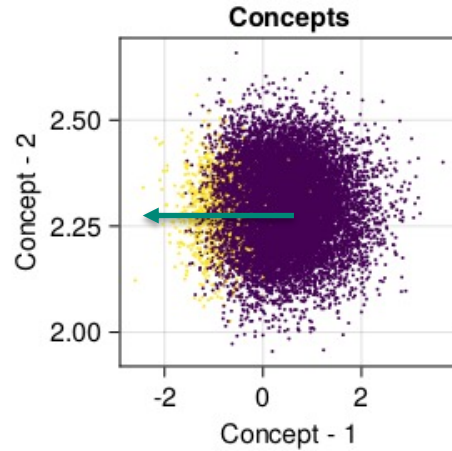
$$\mathcal{L}_y(f(x), y) + \lambda \mathcal{L}_\theta(f(x)) + \eta \mathcal{L}_h(x, \hat{x})$$

- Hilbert-schmidt independence criterion (HSIC)

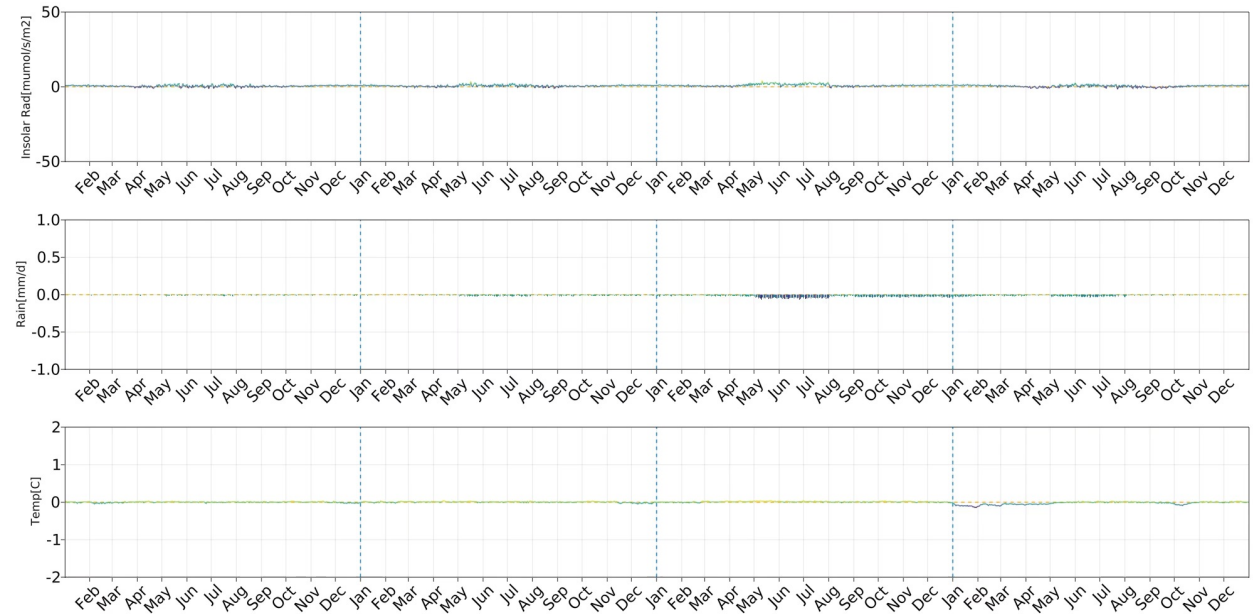
Independent concepts/parameters



Results



- Traversing the concept space
- Difference from the mean climate



Conclusions

- Deep learning can model processes involving long-term memory effects and non-linear driver interactions
- SENN is able to extract explainable concepts responsible for forest mortality
- HSIC loss can disentangle the concept space

References

1. Fatichi, S., Ivanov, V.Y., E. Caporali (2011). Simulation of future climate scenarios with a weather generator, *Advances in Water Resources*, 34, 448–467, doi:10.1016/j.advwatres.2010.12.013
2. Fischer R, Bohn F, Dantas De Paula M, Dislich C, Groeneveld J, Gutiérrez AG, Kazmierczak M, Knapp N, Lehmann S, Paulick S, Pütz S, Rödig E, Taubert F, Köhler P, Huth A (2016) Lessons learned from applying a forest gap model to understand ecosystem and carbon dynamics of complex tropical forests. *Ecological Modelling*, Volume 326, 24 April 2016, Pages 124-133, ISSN 0304-3800, <http://dx.doi.org/10.1016/j.ecolmodel.2015.11.018>.
3. Alvarez-Melis, D., & Jaakkola, T. S. (2018). Towards Robust Interpretability with Self-Explaining Neural Networks. *Advances in Neural Information Processing Systems, 2018-December*, 7775–7784. <https://doi.org/10.48550/arxiv.1806.07538>