

Talk Abstracts

1. On the use of generalized Hoeffding decomposition for data and model explainability (and beyond) by Nicolas Bousquet

Abstract: When considering a phenomenon $Y = f(X)$, where $X = (X_1, \dots, X_d)$, and where f is either unknown or too complex to be directly understood, one seeks to provide local and global interpretations of the behaviour of Y as a function of the components of X . In particular, one typically aims to identify which components have the greatest influence on Y , which interactions among the dimensions of X are the most significant within f , and what role is played by possible correlations among these dimensions.

The Hoeffding decomposition, which notably provides a functional generalization of analysis of variance through functional ANOVA, is a valuable tool for this purpose. Assuming that X is random and that its components may be correlated, it yields representations of f that are unique under suitable assumptions and take the form of additive decompositions into functional terms forming hierarchically orthogonal Riesz bases. These bases can be constructed from orthogonal bases in L_2 spaces, which requires Y to belong to such a space.

These representations generalize the classical functional ANOVA decomposition, which notably leads to the construction of Sobol' indices. Recent results show that such decompositions can be efficiently approximated for both numerical and categorical variables, either when the distribution of X is known or in supervised learning settings where only a sample of input-output pairs (x, y) is available.

Although their approximation raises technical difficulties similar to those encountered in compressed sensing, this approach makes it possible to define refined importance measures—including marginal and interaction-based measures—at both local and global levels. These measures possess useful properties, such as the exclusion of non-causal input dimensions. Beyond interpretability, the Hoeffding decomposition also appears to be a promising tool for improving the construction of surrogate models for complex systems.

2. Free Probability for Neural Network Jacobians, and the Inversion of S-Transforms by Reda Chhaibi

Abstract: In the infinite-width limit, the input-output Jacobian of a feed-forward network at initialization is a product of asymptotically free random matrices, so its law of squared singular values is a free multiplicative convolution of layerwise laws. We extend this beyond constant width by introducing a rectangular multiplicative free convolution — associativity forces the dimension ratio to be carried along as a central extension — and obtain a product formula for the S-transform of the limiting law. Recovering the density then amounts to inverting a multi-valued holomorphic map, an operation long held to be impractical. We give a homotopy scheme: starting high in the upper half-plane where the solution is known, we chain Kantorovich-certified basins of attraction by doubling and dichotomy, with guaranteed convergence to the correct branch and an order of magnitude gained over generic root finding. On 200 random MLP architectures across three datasets, upper quantiles of the predicted spectrum correlate with post-training test accuracy up to $R = 0.85$, at under a minute of CPU each. Some hyperbolicity of the Jacobian, rather than the edge of chaos, emerges as the desirable feature. Joint work with Tariq Daouda and Ez'echiel Kahn, published at NeurIPS 2022 (arXiv:2111.00841).

- 3. Why Transformers Struggle with Distribution-Independent In-Context Learning?** by Nicholas Asher
- 4. A martingale approach for the elephant random walk with stops** by Bernard Bercu
- 5. Introduction to Reinforcement Learning with Human Feedback** by Emmanuelle Claeys
- 6. Stochastic optimization beyond the gradient case** by Gersende Fort
- 7. Sensitivity analysis a quick survey on recent results** by Fabrice Gamboa
- 8. Bias lies also in the learning step of AI algorithms, not only in the data?** by Jean-Michel Loubes
- 9. Document Image analysis: Toward document understanding** by Thierry Paquet
- 10. Causal inference from longitudinal data: Latent** by Myriam Tami