



PhD position in artificial intelligence and accelerator physics

Explainable state observers and interpretable AI for superconducting accelerators and radioactive isotopes identification

Description:

Particle accelerators and radioactive beam production facilities are increasingly complex systems, producing massive volumes of heterogeneous time-series and spectral data. Two strategic GANIL facilities illustrate complementary challenges:

- **SPIRAL2**, with its 26 superconducting cavities, faces instabilities (quenches, multipacting, microphonics) that reduce machine reliability and beam availability. Current anomaly monitoring pipelines (e.g. Time2Feat, tsfresh) have demonstrated clustering of failures but still lack causal interpretability and predictive capability [Tennant et al., 2020; Boukela et al., 2024].
- **SPIRAL1**, producing radioactive ion beams, requires robust isotope identification through noisy β - γ spectroscopy. Classical analysis fails under high background, beam heating, and overlapping decay lines [Kamuda et al., 2020].

Across engineering domains, state observers are a cornerstone of fault detection and diagnosis [Frank, 1990; Chen & Patton, 1999; Isermann, 2006; Ding, 2008]. They reconstruct hidden states and provide residuals that quantify discrepancies between predicted and observed behaviors. These residuals underpin anomaly detection (statistical thresholds, clustering, supervised or self-supervised classification). Observers also act as predictors, enabling early fault alarms and prognostics.

Modern research converges on hybrid pipelines: physics-based observers generate interpretable residuals, while data-driven ML models perform anomaly isolation and explainability [Edelen et al., 2020].

This PhD project positions GANIL at the forefront of this convergence. It will build observer-based interpretable AI frameworks for two critical domains:

1. SPIRAL2 – interpretable anomaly detection and classification of cavity failures.
2. SPIRAL1 – observer-assisted, causal isotope identification under noise.

By uniting observer theory, anomaly detection, prediction, and explainability, this research will create generalizable methods for accelerator operations and nuclear spectroscopy, aligned with international initiatives.

Explainability dimension:

This PhD will embed explainability systematically into the observer+AI framework:

- Observer-based interpretability: residuals are inherently explainable as prediction errors relative to physics-based models [Isermann, 2006].
- Feature attribution: SHAP [Lundberg & Lee, 2017] and LIME [Ribeiro et al., 2016] will be applied to classifiers to highlight the most influential RF or spectral variables.
- Causal discovery: cause–effect graphs will be built from accelerator and spectroscopy data, extending recent surveys [Assaad et al., 2022].
- Uncertainty-aware explainability: predictions will be coupled with calibrated confidence scores (e.g. conformal prediction) to avoid overconfident false alarms.
- Human-in-the-loop trust: explanations will be tested with domain experts (RF operators, nuclear physicists) to ensure usability and adoption.



Expected skills:

Applied physics, control/observer theory, machine learning, time-series/spectral analysis.

Contact: Adnan Ghribi

Phone: +33 (0)2 31 45 46 80

mail: adnan.ghribi@ganil.fr

GANIL, BP 55027, 14076 Caen France

1. Isermann, R. (2006). Fault-Diagnosis Systems. Springer.
2. Chen, J., & Patton, R. (1999). Robust Model-Based Fault Diagnosis for Dynamic Systems. Springer.
3. Ding, S.X. (2008). Model-based Fault Diagnosis Techniques. Springer.
4. Frank, P.M. (1990). Fault diagnosis in dynamic systems using analytical and knowledge-based redundancy. Automatica, 26(3), 459–474.
5. Tennant, C., Carpenter, A., Powers, T., et al. (2020). Superconducting RF cavity fault classification using machine learning at Jefferson Laboratory. Phys. Rev. Accel. Beams, 23, 114601.
6. Boukela, L., Eichler, A., Branlard, J., & Jomhari, N.Z. (2024). Machine learning-aided quench identification at the European XFEL. arXiv:2407.08408.
7. Edelen, A.L., Neveu, N., Frey, M., et al. (2020). Machine learning for orders-of-magnitude speedup in multiobjective optimization of accelerator systems. Phys. Rev. Accel. Beams, 23, 044601.
8. Assaad, C.K., Devijver, E., & Gaussier, E. (2022). Causal discovery methods for time series. JAIR, 73, 767–819.
9. Kamuda, M., Zhao, J., & Huff, K. (2020). Semi-supervised isotope identification. NIM-A, 954.
10. Ribeiro, M.T., Singh, S., & Guestrin, C. (2016). "Why Should I Trust You?" Explaining the Predictions of Any Classifier. KDD 2016.
11. Lundberg, S.M., & Lee, S.I. (2017). A unified approach to interpreting model predictions (SHAP). NeurIPS 2017.
12. Qiu, C., Pfrommer, T., Kloft, M., Mandt, S., & Rudolph, M. (2021). Neural transformation learning for deep anomaly detection beyond images. ICML 2021.
13. Reiss, T., & Hoshen, Y. (2023). Mean-shifted contrastive loss for anomaly detection. AAAI Conf. AI, 37, 2155–2162.