

AI/ML at FACET-II
(E331 / E341)

Broad goals for AIML program

- **Enable new capabilities in beam customization and delivery**
 - Advanced phase space manipulation for single and two bunch
 - High-dimensional tuning across multiple objectives
 - e.g. transverse and longitudinal phase space, minimize beam losses, improve beam stability
 - Automated injector startup, charge change, emittance growth minimization
 - Ability to tune complex components: e.g. sextupole tuning + w-chicane
- **Enable new capabilities in beam measurement / prediction**
 - 6D phase space reconstruction (GPSR)
 - ML-enhanced virtual accelerator → continual learning
- **Work with other experiments/ops to transition above to regular use**

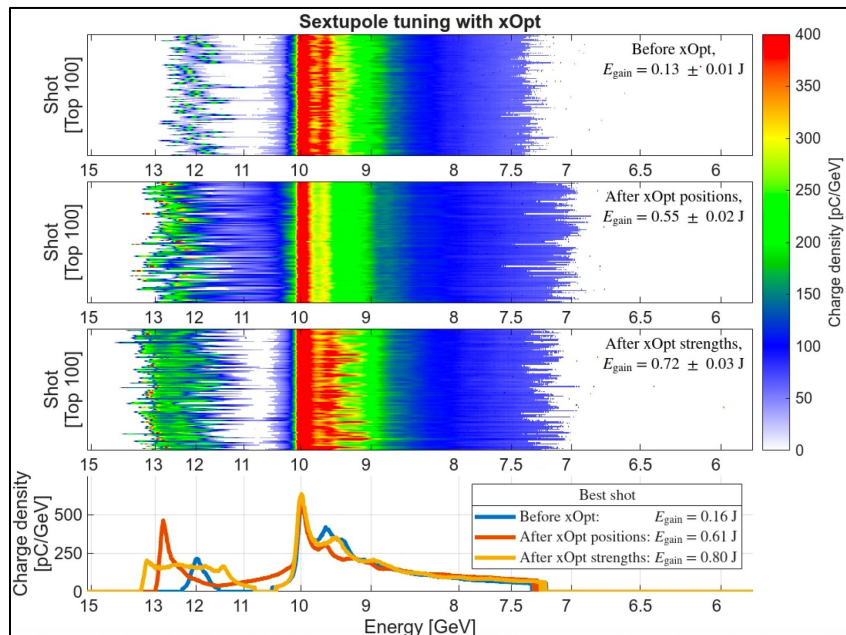
Relation to broader AIML activities

- Have several AI/ML research grants:
 - Hardware-Aware AI (HAAI) focused on system modeling and two bunch tuning
 - Continual learning, novel representations for system modeling
 - Dynamic two-bunch tuning (time/energy separation)
 - Genesis MOAT (multi-office accelerator team) - agentic AI and digital twins for autonomous accelerators
- SLAC ISDCI well positions for Genesis AmS + HPC in the loop with experiments (including FACET-II)
- Effort to pitch FACET-II as a fully AIML integrated test bed
 - FACET-II experiment time has been heavily used to demonstrate AIML algorithms and software

S20 Tuning on sextupoles (E331 → E300 success story)

Initial work on
sextupole mover
optimization with
Xopt in E331

Now transferred
over to E300 with
excellent results
and substantial
further
development by
E300 folks



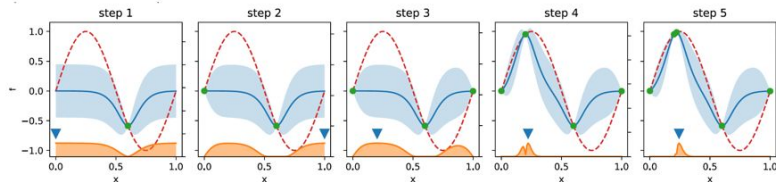
**5x overall
improvement in
energy gain for
plasma-based
acceleration
with ML-driven
sextupole tuning**

S. Rego

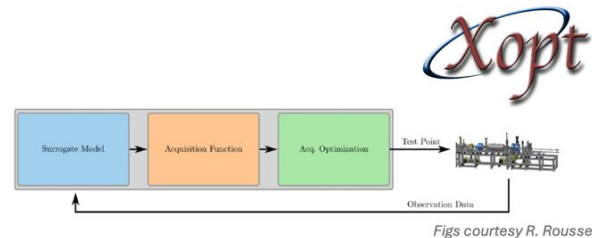
D. Storey and E300 team

From 2024 initial tests

Bayesian Optimization of Sextupole Mover Settings

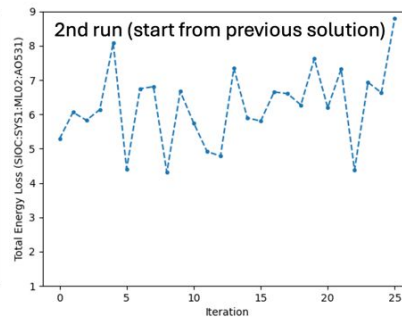
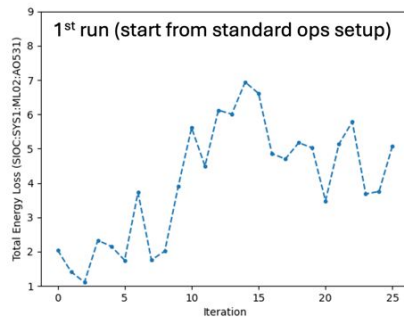


Bayesian Optimization learns a Gaussian Process model on-the-fly during sampling. Model predictions and uncertainties guide the next point to sample.
→ no need for previous data
→ easy-to-adapt implementations of advanced BO techniques available in Xopt



Figs courtesy R. Roussel

More BO info: <https://arxiv.org/abs/2312.05667>
Xopt: <https://github.com/ChristopherMayes/Xopt>



Ran separate optimizations starting from initial operations setup

8 variables – 4 sextupole magnet mover positions in x and y

Objectives are beam size on screen or plasma parameters (measured on images)

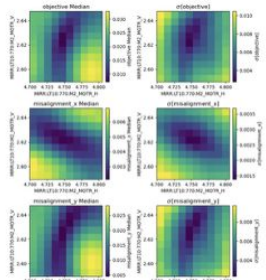
Plan to put into Badger GUI for general use in operations



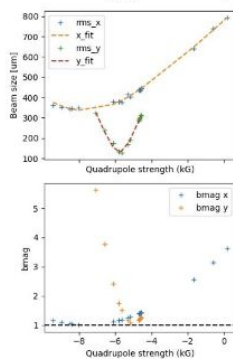
Moving toward linking together chains of tasks autonomously

Autonomous configuration of the FACET-II injector beamline w/o human intervention in < 14 min (standard is ~1hr +)

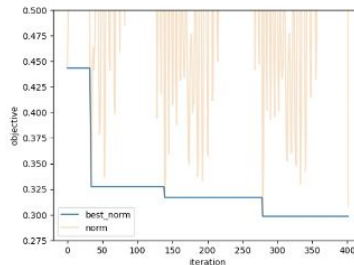
Laser alignment onto cathode (1 min)



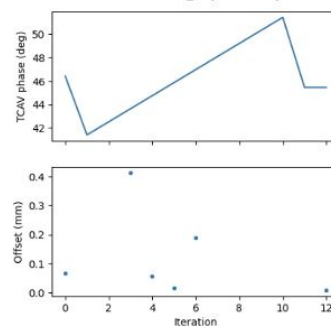
Emittance optimization + matching (7.5 min)



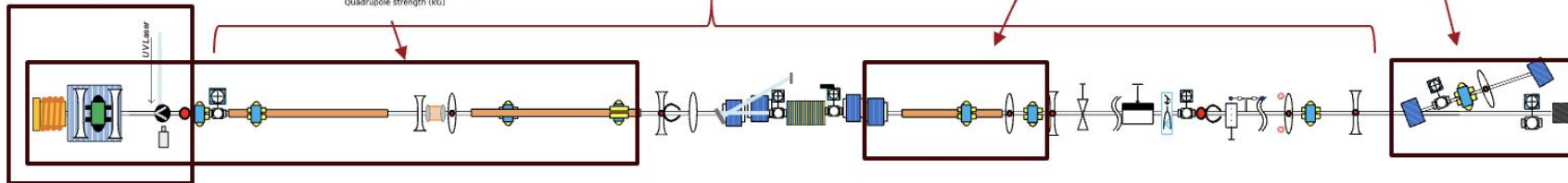
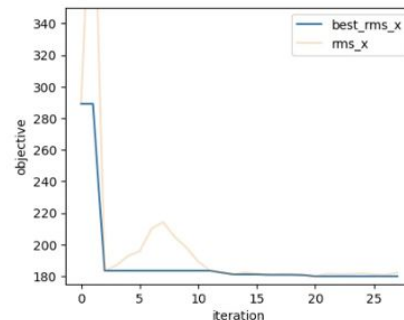
Beam steering (3 min)



TCAV Phasing (7 sec)



Energy spread minimization (2 min)

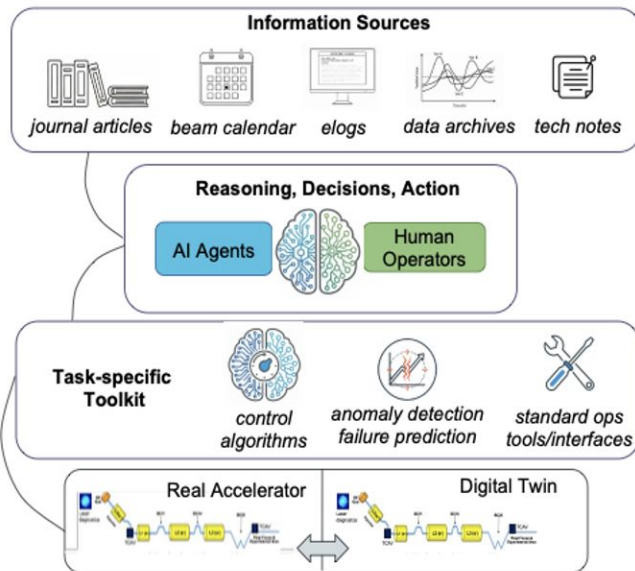


Agents Lower Barriers in Advanced Tool Use



Agentic AI (Osprey*) being developed for complex operations and design workflows for accelerators

Osprey orchestrates tool use and accesses disparate accelerator information sources



Per-knob bounds (updated)

Knob	Current (live)	Box limits	HW limits	Proposed Range (LL-H)
QAA0-1700-000-000-000-000	-48.32	(none)	[-132, 332]	[-42.52, -37.702]
QAA0-1700-000-000-000-000	47.03	(none)	[-132, 332]	[45.23, 50.48]
QAA0-1700-000-000-000-000	-45.54	(none)	[-132, 332]	[-47.04, -43.342]
QAA0-1700-000-000-000-000	47.47	(none)	[-132, 332]	[45.27, 50.47]

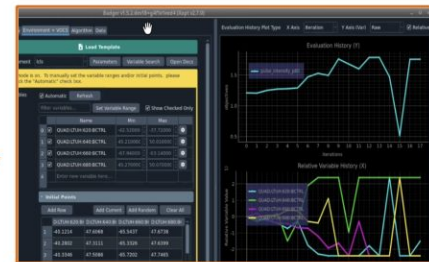
All within hardware limits [X] All current values centered in proposed ranges [X]

Step 5 - Proposed Routine

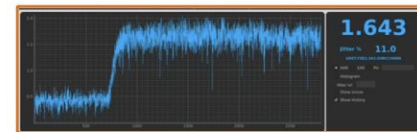
Field	Value
Name	H8R FEL tune-up L70m watching quads v2
Box	130s (5s/hr)
Knobs	4 L70m watching quads (32.4 around current)
Objective	quads.intensity_pos → MAXIMIZE
Algorithm	Expected Improvement (128 MC samples, LBFGS, 24 restarts)
Initial points	1 current + 3 random (30s FILL)
Based on	Apr 15 best run + updated live machine state

user powered Claude's questions:
 - How to save this updated routine? The ranges are re-centered on the current machine state (which has shifted from the previous optimization). → Save as proposed
 - Calling badger-activate. (confirm to export)

New optimization routine



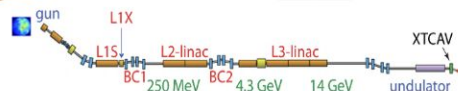
Automatic execution on machine via Badger tuning tool



Rapid improvement in FEL performance

Prompt to completion: ~ 10 minutes

natural language description of tuning task



Have demonstrated in live accelerator tuning at SLAC → agent uses complex AIML optimization tools, reducing burden on operators to configure them

Agent-driven startup for the injector

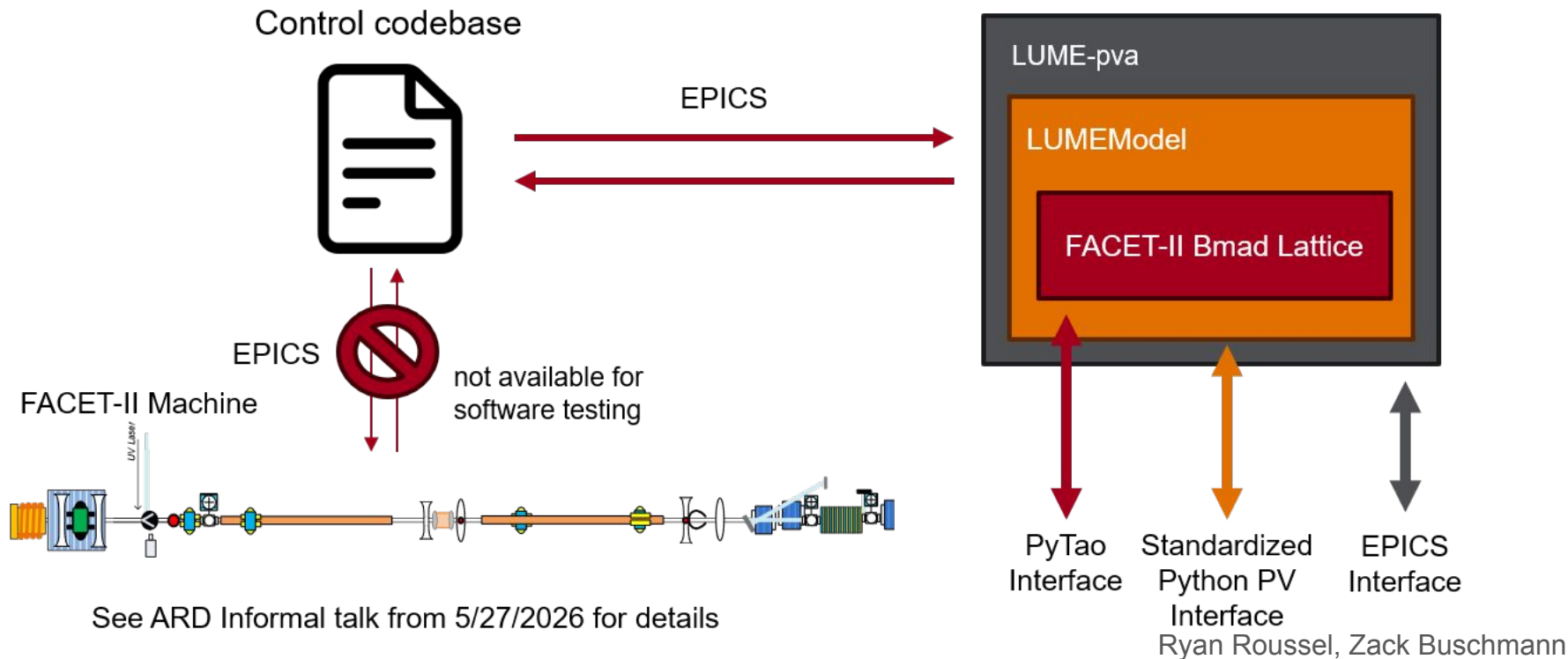
Osprey/Otter → used Xopt tools from each step given high-level task / instructions for injector startup

Ran several charge changes and tuning: 1.6 nC, 2 nC, 2.2 nC

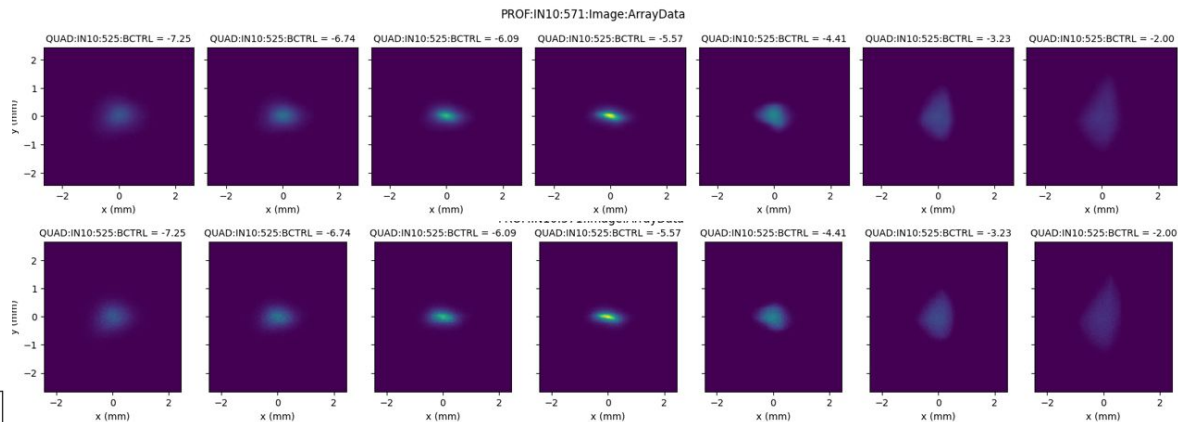
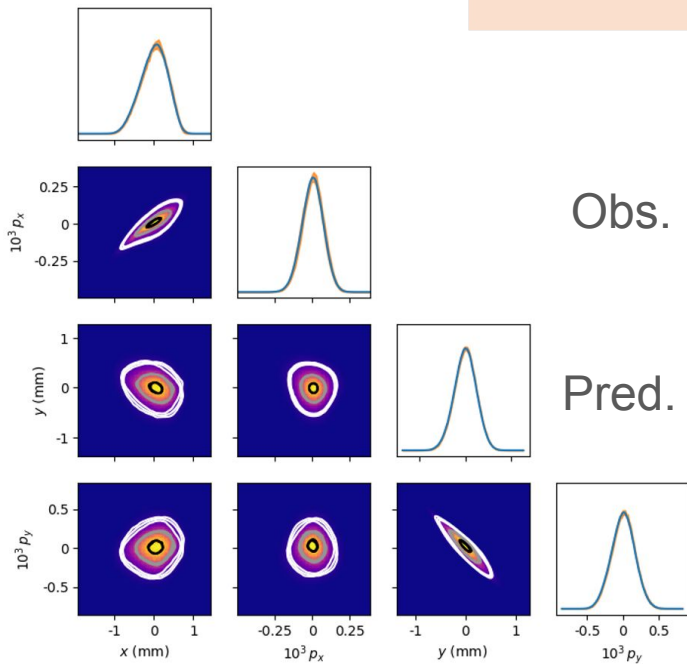
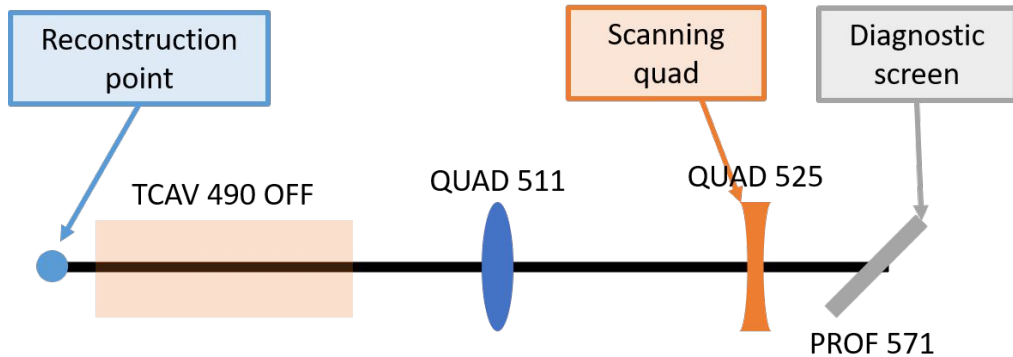
- ~20 minutes for entire task chain
- Recovered roughly equivalent or better emittance in each case (although back to 1.6 nC had poorer match)

Virtual Accelerator

Prior to deployment on the machine, algorithms were tested on a virtual accelerator (VA) of FACET-II using LUME



GPSR



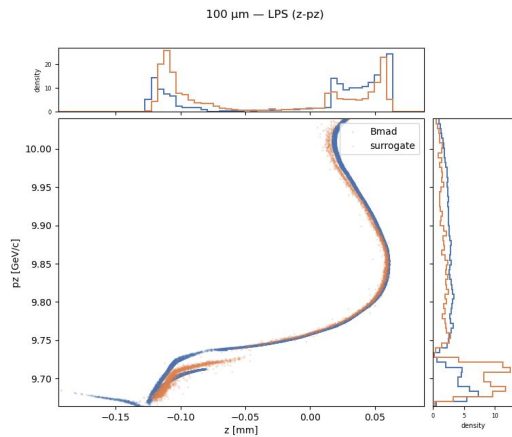
Uses LUME-based FACET virtual accelerator to reconstruct distribution

Dynamic two-bunch tuning

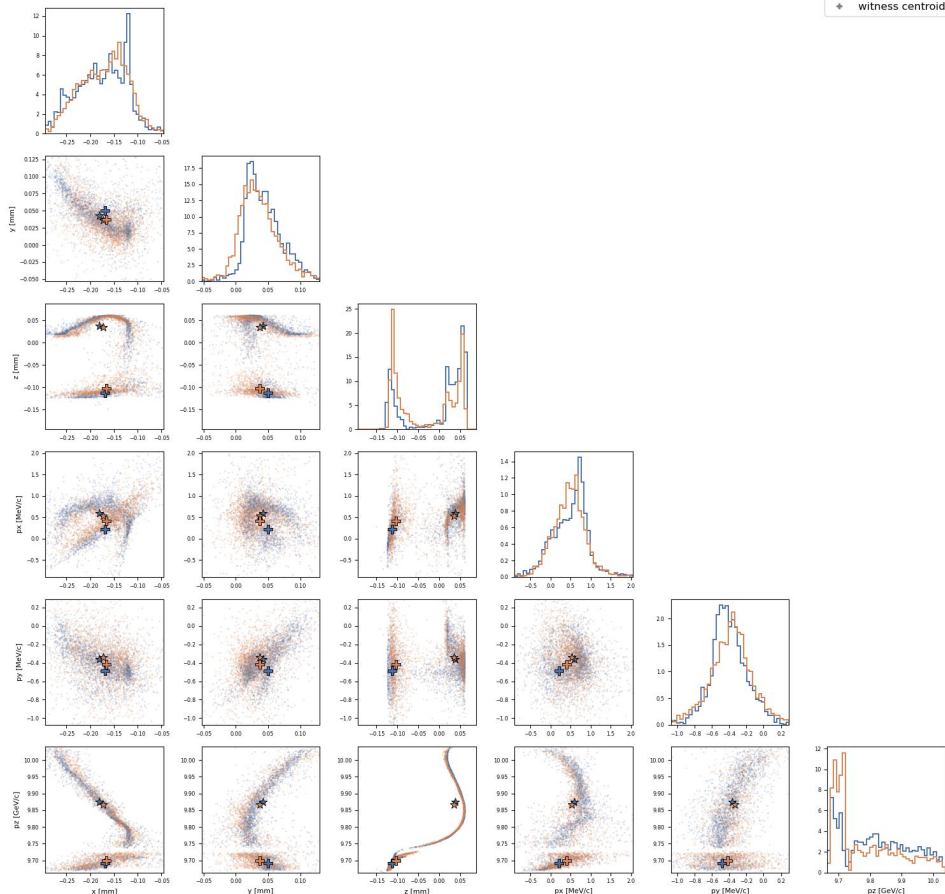
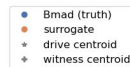
Goal: specify timing and energy separation and be able to slew across these easily

Model-based reinforcement learning approach

Initial results from Bmad simulation



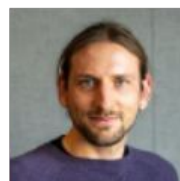
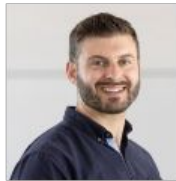
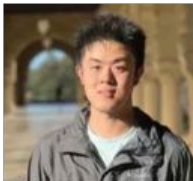
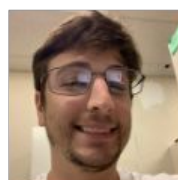
100 μm setpoint — Bmad vs surrogate



Papers we are aiming for

1. Joint with E300 – ML based sextupole tuning and outcomes for plasma performance
2. Agent-driven tuning and/or autonomous injector startup → plan to start using again during machine startup
3. Dynamic two-bunch tuning with reinforcement learning (slew time/energy separation) → will request beam time when machine starts up
4. Model calibration / sim2real for FACET-II

Thanks to many contributors (not exhaustive)



Multi-office Accelerator Team (MOAT) - ModCon seed

- **Brings together DOE BES, HEP and NP (now HENP)**
 - 7 national labs: LBNL+ANL+BNL+FNAL+JLAB+ORNL+SLAC
- **MOAT vision: treat nation's accelerator ecosystem as a single, intelligent system**
 - using AI to connect expertise across labs and offices
 - views large accelerators as complex, dynamic socio-technical systems
 - *distributed teams of teams, thousands of physical components, vast stores of scientific knowledge, rigorous safety and management frameworks*
- **The MOAT-seed project builds initial foundations toward vision:**
 - deployment of shared accelerator agentic AI software across facilities (8 to date)
 - deployment of digital twin prototypes across facilities (3 to date)
 - AI-assisted accelerator designs
- **Genesis Mission's AmSC+ModCon provide the interconnecting tissue of computing + AI capabilities**
- **SLAC is a major leader/contributor within MOAT:** due to prior investment / infrastructure, experienced teams, and availability to test on accelerators
 - AI/ML methods development and shared software development
 - Demonstrations across different SLAC accelerator use-cases (*aided by S3DF/ISDCI and long-standing SLAC leadership/investment in AI/ML*)

