

HEPTAPOD: Orchestrating High Energy Physics Workflows Towards Autonomous Agency

Tony Menzo^{1,2,*}, Alexander Roman[†], Sergei Gleyzer, Konstantin Matchev

¹*Department of Physics and Astronomy, University of Alabama, Tuscaloosa, AL 35487, USA*

George T. Fleming, Stefan Höche, Stephen Mrenna, Prasanth Shyamsundar

²*Fermi National Accelerator Laboratory, Batavia, IL 60510, USA*

December 12, 2025

Abstract

Many theoretical and experimental workflows in high-energy-physics (HEP) stand to benefit from recent advances in transformer-based large language models (LLMs). While early applications of LLMs focused on text generation and code completion, modern LLMs now support *orchestrated agency*: the coordinated execution of complex, multi-step tasks through tool use, structured context, and iterative reasoning. We introduce the **HEP Toolkit for Agentic Planning, Orchestration, and Deployment** (HEPTAPOD), an orchestration framework designed to integrate LLMs into general HEP workflows spanning theoretical calculations, simulation, and data analysis. The framework enables LLMs to interface with domain-specific tools and to construct and manage diverse HEP pipelines while preserving transparency, reproducibility, and human oversight. To demonstrate these capabilities, we present a representative case study in the context of a Beyond the Standard Model (BSM) Monte Carlo signal validation that spans model generation, event simulation, and analysis within an established, reproducible workflow. HEPTAPOD provides a structured and auditable layer between human researchers, LLMs, and computational infrastructure, establishing a foundation for human-in-the-loop, agent-assisted workflows across high-energy physics.

*Co-first author, corresponding author: amenzo@ua.edu

[†]Co-first author.

Contents

1	Introduction	3
2	An example workflow	4
3	Orchestration architecture and philosophy	6
3.1	Schema-validated tools	6
3.2	LLM-compatible event data structures	9
3.3	Run-card-driven orchestration	10
3.4	Context propagation and stateful execution	10
3.5	Orchestration versus scripting	11
4	Agent-orchestrated workflow	13
4.1	Example conversation	14
5	Conclusions	27
A	Tools	29
A.1	Event Generation	29
A.2	Data Conversion	31
A.3	Jet Clustering	31
A.4	Kinematic Analysis	31
A.5	Event Selection	32
A.6	Advanced Analysis	32

1 Introduction

Large language models (LLMs) have evolved from text generators into general-purpose computational interfaces, capable of planning and executing multi-step workflows through structured prompting and function calls [1–6]. Their capabilities appear to scale predictably with model size, dataset quality, and training compute, following empirical scaling laws that remain robust across architectures and domains [7, 8]. This progression has introduced the notion of *LLM agency*, in which models operate autonomously or cooperatively to complete complex objectives [9–11]. As investments in large-scale compute infrastructure continue, these systems are expected to improve in reasoning, tool use, and generalization, suggesting a near-term transition from passive language processors to reliable orchestrators of scientific workflows [12].

High-energy physics (HEP) presents a natural testing ground for these capabilities. Modern theory and phenomenology rely on sophisticated computational stacks spanning symbolic algebra, matrix-element generation, parton showering and hadronization, detector simulation, and numerical analysis, all of which must operate in coordinated, reproducible sequences [13]. Yet most pipelines remain manually managed, requiring researchers to orchestrate long chains of interdependent calculations, propagate parameters across tools, interpret heterogeneous outputs, and debug failures at any stage.

Concurrently, machine learning (ML) methods have reshaped HEP research across simulation [14] (including phase-space integration and event unweighting [15, 16], hadronization modeling [17–21], and fast detector simulation using generative models [22]), theory inputs such as parton distribution function determinations [23, 24], event reconstruction [25], triggering [26, 27], anomaly detection [28, 29], unfolding [30, 31], and analysis [32, 33]. Transformers in particular have become standard architectures for particle-cloud reasoning in jet physics, with applications ranging from task-specific jet classification [34, 35] to autoregressive generative modeling of jet radiation [36] and more recently to pretrained foundation models for jet representations [37, 38]. Related but more general efforts have begun to explore pretrained transformer foundation models for collider-level data [39]. Early applications of transformer-based LLMs in HEP have focused on code generation, documentation, and integration with existing software infrastructures [40, 41]. More recently, agentic frameworks have begun coordinating multi-tool pipelines; for example, **ArgoLOOM** provides an agentic AI framework for cosmology, collider physics and nuclear science and produces reproducible artifacts such as run cards and plots [42].

We introduce the **HEP Toolkit for Agentic Planning, Orchestration, and Deployment** (HEPTAPOD)¹, an orchestration framework that integrates LLMs directly into HEP theory pipelines. The agentic execution layer is implemented using the ORCHESTRAL AI framework [44].² Rather than pursuing full autonomy, HEPTAPOD emphasizes human-in-the-loop control: the LLM acts as a workflow coordinator that invokes domain-specific tools, interprets intermediate outputs, and proposes next steps, while researchers retain oversight and validation. This provides the necessary

¹Not to be confused with the image autoregressive model for language modeling on visual signals introduced in [43]. That model, also referred to as HEPTAPOD, is named after the alien species in the film *Arrival* (2016).

²A companion paper showcases an analogous application of ORCHESTRAL AI for the analysis and simulation of exoplanet transit spectroscopy [45].

structure so that increasingly-capable scientific agents can operate safely and reproducibly within established computational workflows.

The HEPTAPOD framework is well aligned with emerging initiatives that seek to leverage agentic AI systems for discovery-driven science.³ By abstracting theory tasks into modular, schema-validated tools, HEPTAPOD enables applications ranging from effective-field-theory construction to automated event-generation pipelines and symbolic-numerical translation. More broadly, it provides a blueprint for how agentic LLMs can augment scientific reasoning without displacing human judgment, establishing a foundation for reproducible, adaptive, and interpretable AI-assisted research in theoretical physics.

The rest of this paper is organized as follows. Section 2 introduces a representative leptoquark workflow that will serve as a running example throughout the paper. Section 3 describes the orchestration architecture underlying HEPTAPOD, including schema-validated tools, LLM-friendly event formats, run-card-driven coordination, and the distinction between orchestration and traditional scripting approaches. Section 4 demonstrates the complete framework through a BSM leptoquark parameter scan and signal validation. Section 5 concludes with a discussion of future work and perspectives on increasingly agentic scientific systems.

2 An example workflow

To motivate the design goals of HEPTAPOD and to illustrate the challenges involved in coordinating multi-stage high-energy-physics workflows and calculations, we introduce a representative example that will serve as a running case study. Specifically, we consider the Monte Carlo simulation pipeline required to study new BSM models that are not available by default in general-purpose event generators. In this situation, model implementation, event generation, and downstream simulation must be manually composed across multiple software packages. The design and automation of this pipeline was the subject of a series of international workshops, Monte Carlo Tools for Beyond the Standard Model Physics (MC4BSM), which ran in the 2000’s and 2010’s [47]. The example considered here is very close to the test case study demonstrated in the MC4BSM tutorials [13].

For concreteness, we consider a scalar leptoquark [48]

$$S_1 \sim (\bar{\mathbf{3}}, \mathbf{1}, 1/3), \quad (2.1)$$

with a simplified interaction Lagrangian

$$\mathcal{L} \supset (D_\mu S_1)^\dagger (D^\mu S_1) + y_{ij} \overline{u_{Ri}^c} e_{Rj} S_1 + \text{h.c.} \quad (2.2)$$

Assuming diagonal Yukawa couplings, pair production proceeds dominantly through QCD,

$$pp \rightarrow S_1 S_1^\dagger \rightarrow (\ell^+ j)(\ell^- j), \quad (2.3)$$

yielding a characteristic $2\ell + 2j$ topology. The task is to perform MC signal generation over a scan

³In particular, it supports the objectives of the recently launched Genesis Mission, including the Discovery Science challenge of “Understanding the universe, from quarks to cosmos” [46].

of the leptoquark mass m_{S_1} and predict the distribution of the reconstructed leptoquark mass $m_{\text{LQ}}^{\min}$, which for concreteness is taken as the smaller of the two reconstructed masses of the leptoquark candidates $m_{\text{LQ}}^{(1)}$ and $m_{\text{LQ}}^{(2)}$

$$m_{\text{LQ}}^{\min} = \min \left\{ m_{\text{LQ}}^{(1)}, m_{\text{LQ}}^{(2)} \right\}, \quad (2.4)$$

where the lepton-jet assignments are selected to minimize the absolute difference $|m_{\text{LQ}}^{(1)} - m_{\text{LQ}}^{(2)}|$. A typical phenomenological study varies m_{S_1} over a set of benchmark values, e.g., (1.0, 1.5, 2.0 TeV), generating and analyzing events for each mass point. This requires coordinated propagation of model parameters, consistent run-card configuration across tools, and synchronized analysis logic across the scan. The scan is simple, yet has the essential structure of a multi-stage HEP workflow: configuration must propagate consistently across several tools, intermediate files must be handled carefully, and analysis logic must remain synchronized across scan points.

Although not included here, realistic validation pipelines must also include the dominant SM backgrounds to this signature. For the $(\ell j)(\ell j)$ final state, the most relevant backgrounds are

- Drell-Yan + jets: $pp \rightarrow Z/\gamma^* + \text{jets} \rightarrow \ell^+ \ell^- + \text{jets}$,
- Top-quark pair production: $pp \rightarrow t\bar{t} \rightarrow (b\ell^+\nu)(\bar{b}\ell^-\bar{\nu})$,
- Diboson production: WW , WZ , or ZZ events with associated jets.

Each background requires its own run card, generation chain, and analysis flow, producing multiple parallel branches that must remain consistent. These branches are straightforward to implement within the orchestration framework because they reuse the same tool chain with different run-card inputs. Here we focus solely on the signal to illustrate the orchestration framework.

For this example, a complete HEP workflow can be instantiated in several ways. One representative realization is considered here (see fig. 1), though comparable pipelines using, for example, SHERPA [49] follow the same basic structure:

1. **Model generation:** convert the FEYNRULES [50] model file into a Universal FeynRules Output (UFO) directory.
2. **Parton-level event generation:** use MADGRAPH5_AMC@NLO [51] to generate events for each mass point specified in the run card via a user-defined parameter scan.
3. **Showering and hadronization:** use PYTHIA [52] to process the LHE events into particle-level final states.
4. **Jet reconstruction:** apply anti- k_T or Cambridge/Aachen clustering via FASTJET [53].
5. **Resonance reconstruction:** pair leptons and jets, compute $m_{\text{LQ}}^{(1)}, m_{\text{LQ}}^{(2)}$, and derive $m_{\text{LQ}}^{\min}$.

Careful management of run cards, file paths, tool configurations, and intermediate data formats is necessary to execute this sequence reliably. Manual pipelines are time-intensive and error-prone. In later sections we return to this example and show how HEPTAPOD treats each phase as a schema-validated tool, propagates scan metadata through the workflow automatically, and uses run cards as explicit coordination boundaries. This allows an LLM to orchestrate the full analysis while preserving transparency, reproducibility, and human oversight.

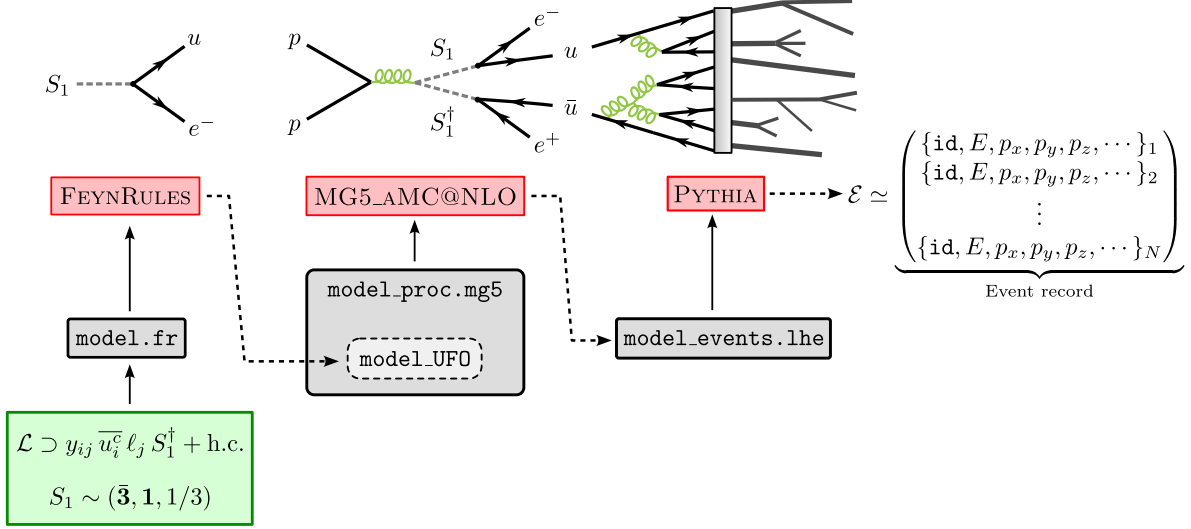


Figure 1: One example of a multi-phase software workflow for a BSM leptoquark study. A `model.fr` file is implemented in FeynRules and exported as a UFO directory (`model_UFO/`). The UFO is consumed by MG5_AMC@NLO, which loads the process card `model_proc.mg5` and generates parton-level events written to `model_events.lhe`. These events are then passed to PYTHIA for showering and hadronization, producing an event record $\mathcal{E}$ containing particle identifiers and four-momenta $\{p_x, p_y, p_z, E, \text{id}, \dots\}$.

3 Orchestration architecture and philosophy

The central objective of HEPTAPOD is to facilitate and enable large language models (LLMs) to reliably coordinate high-energy-physics (HEP) workflows while preserving transparency, reproducibility, and human oversight. The leptoquark workflow introduced in section 2 highlights the central challenge in automating HEP analyses: while each individual phase of the pipeline is well supported by mature software packages, the coordination of these components remains largely manual. Researchers must construct run cards, manage intermediate files, ensure parameter consistency, and sequence tool execution correctly across multiple stages. These tasks are labor-intensive and error-prone, especially when extended across parameter scans or repeated for large model spaces. The following subsections outline the architectural features of HEPTAPOD that address these coordination challenges. A high-level view of the HEPTAPOD architecture and its interaction between agents, tools, and external HEP software is shown in fig. 2.

3.1 Schema-validated tools

Each scientific capability in HEPTAPOD is exposed to the agent as a *tool*: a Python class with a structured JSON schema defining its inputs and outputs. Tools serve as the fundamental interface between the agent and external HEP software. Within the orchestration framework, tools play three roles: 1. to formalize how the agent invokes domain-specific capabilities, 2. to return structured outputs that downstream steps can consume reliably, and 3. to define clear execution boundaries that support reproducibility and human oversight.

Tool schemas distinguish between *runtime fields*, which the agent provides when invoking the tool, and *state fields*, which are injected by the orchestration layer and not visible to the agent. Runtime fields correspond to the scientific and programmatic degrees of freedom, like relative paths, algorithmic configurations, and parameter choices. State fields supply environmental information such as workspace locations, internal defaults (*e.g.* dependency paths), or metadata from previous phases. This separation keeps the agent’s reasoning at the physics level while insulating it from unnecessary overhead from the details of the filesystem or workspace (*a.k.a.* sandbox).

A tool’s schema is defined directly by its class annotations, from which the orchestration layer derives validation rules that enforce argument conventions and ensure that all values exchanged between the agent and the execution environment are checked prior to tool execution. For example, a simplified PYTHIA interface might declare:

```
class PythiaFromRunCardTool(BaseTool):
    """
    Generate hadron-level events using Pythia8 from a .cmnd card.

    Inputs (runtime):
        - command_card: path to Pythia8 .cmnd configuration file
        - nevents: number of events to generate
        - seed: optional random seed
        ...

    Behavior:
        1. Initialize Pythia8 with the provided run card
        2. Generate nevents events with optional fixed seed
        3. Write structured event records to JSONL format
        ...

    Output (JSON):
        {"status": "ok", "events_jsonl": "<path>", ...}
    """
    # Runtime fields
    command_card: str = RuntimeField(
        description="Path to Pythia8 .cmnd configuration file")
    nevents: int = RuntimeField(
        description="Number of events to generate")
    seed: Optional[int] = RuntimeField(
        default=None, description="Random seed (optional)")
```

An agent invoking this tool submits a structured tool call, which is validated before execution. Schema validation prevents failures from those cases when runtime inputs do not satisfy the tool’s declared interface. When failures do arise from external software (*e.g.* FEYNRULES, MADGRAPH,

PYTHIA), tools return structured error objects that the agent can interpret and act upon. This provides predictable failure modes in place of unstructured logs.

As a design convention, tools also return structured JSON outputs. This ensures that each tool produces data in a form that downstream steps can consume without ad hoc parsing or heuristics. The output schema specifies the fields and their types, which guarantees that intermediate products such as UFO directories, LHE file paths, jet collections, or analysis summaries have a predictable structure. This removes ambiguity at tool boundaries and provides the agent with machine-checkable data structures rather than free-form text, which is essential for coordinating multi-step scientific workflows. An example tool invocation and the associated JSON output for the `PythiaFromRunCardTool` is shown below:

```
PythiaFromRunCardTool(
    command_card="cards/run_ppbar.cmd",
    nevents=5000,
    seed=12345
)
→ {"status": "ok", "events_jsonl": "data/run_01/events.jsonl", ...}
```

Crucially, in addition to the schema itself, the *tool docstring* (the Python docstring, a free-form text description attached to a function or class that explains what it does and how to use it) plays a critical role in guiding the agent’s behavior. Each tool class carries a domain-specific docstring that describes its scientific purpose, the meaning of its runtime fields, expected behavior, and the semantics of its outputs. This docstring functions as the tool’s prompt: when the agent considers invoking the tool, the orchestration layer provides both the schema and the docstring, giving the agent a structured interface (the JSON schema) and a semantic interface (the natural-language description). The tool docstring is extracted and included in the tool specification sent with every LLM API call, ensuring that the agent has immediate access to semantic context needed to understand what parameters to provide and how to interpret the returned fields without relying on conversational memory. The combination is essential. The schema tells the agent what inputs are allowed; the docstring tells the agent why and when the tool should be used.

Because outputs are returned as structured JSON, the tool docstring also defines the interpretation of those fields. For example, the MADGRAPH tool docstring specifies that when a parameter scan is detected, the returned `runs` array contains objects with `run_id`, `lhe_file`, `scan_params`, and `cross_section` fields, where all file paths are relative to the sandbox. This coupling between semantic guidance and machine-checkable structure enables the agent to reason about workflows without relying on free-form text generation or fragile pattern matching. The agent uses the docstring to decide *which* tool to invoke and the schema to decide *how* to invoke it.

Together, the schema and the docstring act synergistically: the docstring communicates intent and scientific context; the schema enforces shape, type, and validity. This is ultimately what allows tools to operate as reliable, composable units within multi-step workflows.

3.2 LLM-compatible event data structures

Effective orchestration requires that intermediate physics data, particularly event-level outputs, be stored in a form that is both machine-actionable and directly interpretable by a language model, enabling the agent to participate meaningfully in validation, debugging, and exploratory analysis. For this reason, HEPTAPOD introduces a structured, text-based representation for intermediate event data, based on a line-delimited JSON (JSONL) format referred to throughout as the `evtjsonl` event format.

Traditional HEP event containers and event-record formats such as ROOT [54], HDF5, HEPMC [55] and LHE [56] are optimized for storage density and I/O throughput in large scale simulation and analysis workflows. While well suited to large-scale production and offline analysis, these formats are opaque to language models: they require specialized libraries for deserialization and expose internal structure only indirectly. As a result, an LLM cannot inspect individual events, validate intermediate outputs, or reason about event content without invoking external tools. In contrast, the JSONL representation encodes each event as a self-contained record that can be examined directly within the conversational context. This allows the agent to inspect representative events, verify kinematic properties, and respond to qualitative queries about the data without leaving the orchestration layer.

The `evtjsonl` format standardizes this representation across all workflow stages while remaining intentionally extensible for future use cases (*e.g.* long-lived particle production, *etc.*). Each event record contains an `event_id` field and a `data` object, which holds either a `particles` array for particle-level events or a `jets` array for reconstructed jet collections. Particle objects include kinematic four-vectors (p_x, p_y, p_z, E) , mass m , PDG identifier `id`, and status codes; jet objects carry four-vectors and jet-specific metadata. This uniform structure allows tools to consume event data predictably: all event generators, shower algorithms, jet clustering tools, and analysis selections produce and expect the same format, eliminating the need for ad hoc parsing or format conversions between workflow stages.

The line-delimited structure of JSONL also enables incremental processing. Because each line is an independent record, tools can iterate through events sequentially, processing and writing results one event at a time without loading entire datasets into memory. The format also simplifies parallelization and filtering: events can be distributed across workers or selected on-the-fly using standard command-line tools. This design choice reflects a broader principle within HEPTAPOD: intermediate data formats should be interpretable not only by the orchestration layer and downstream analysis tools, but also by the agent itself.

While JSONL sacrifices storage compactness relative to compressed formats, HEPTAPOD mitigates this through interoperability tools. `LHETOJSONLTool` imports Les Houches events from matrix-element generators; `EventJSONLToNumpyTool` and `JetsJSONLToNumpyTool` export to zero-padded NumPy arrays compatible with vectorized analysis libraries and machine-learning frameworks such as JAX, PyTorch, and TensorFlow. This hybrid approach retains the agent-facing interpretability of JSONL during orchestration while enabling seamless integration with the broader HEP and ML ecosystems at analysis endpoints.

The uniform adoption of `evtjsonl` across the tool suite exemplifies a core design philosophy: data structures should be legible to both machines and agents, enabling the LLM to participate meaningfully in validation, debugging, and exploratory analysis.

3.3 Run-card-driven orchestration

Most major HEP collider simulation frameworks, including FEYNRULES [50], MADGRAPH [57], PYTHIA [52], SHERPA [49], HERWIG [58], DELPHES [59], GEANT4 [60], and RIVET [61], are configured via human-readable run cards or model files. These declarative configuration files specify model parameters, beam conditions, detector settings, event-selection criteria, and analysis instructions. HEPTAPOD adopts run cards as the canonical interface between the agent and external HEP codes because they provide a stable, interpretable, and versionable surface for steering scientific workflows. Specifically, using run cards as the orchestration boundary provides three key advantages:

Reproducibility and transparency. Run cards define complete, standalone configurations that can be version-controlled, archived, and executed independently of the agent. Each decision made by the LLM is therefore captured in a persistent artifact, enabling fully auditable workflows.

Compatibility with RAG and structured prompting. Because run cards follow stable syntax conventions, they are highly amenable to retrieval-augmented generation (RAG). The agent can be contextualized with validated examples, available parameter dictionaries, and relevant excerpts from official manuals, allowing it to generate syntactically correct and meaningful configurations without memorizing the entire documentation. This reduces the likelihood of malformed configurations and improves the reliability of free-form prompting. We leave a full implementation of RAG-based run-card generation to future iterations of the framework.

Human-approval checkpoints. Each generated run card constitutes a natural boundary at which the researcher may inspect, modify, or approve the configuration before execution. This preserves scientific judgment while still benefiting from automated orchestration.

3.4 Context propagation and stateful execution

HEPTAPOD maintains a structured conversation state or *context* that records tool outputs and relevant metadata. Later steps can reference this state directly, rather than relying on the agent to recall filenames or parse textual logs. For example, when a run card includes a user-defined parameter scan (e.g., multiple values of a mass or coupling), the MADGRAPH tool automatically detects the scan structure by inspecting the resulting run directories for characteristic naming patterns (e.g., `run_01`, `run_02`, ...) and parsing the scan summary files generated by MADGRAPH. Instead of returning only LHE file paths, the tool produces structured metadata:

```
{ "scan_detected": true,
  "n_runs": 3,
  "runs": [
    {"run_id": "run_01", "lhe_file": "...", ...},
```

```

{"run_id": "run_02", "lhe_file": "...", ...},
{"run_id": "run_03", "lhe_file": "...", ...}
}]

```

associating each scan point with its corresponding parameter values, file locations, and computed cross sections. This metadata is written into the conversation state, allowing subsequent steps, such as PYTHIA, to access the scan information directly without requiring the agent to interpret directory names or rely on textual recall.

Additionally, tools automatically detect and substitute placeholder values in run cards at execution time. For instance, a MADGRAPH command card may contain `import model [[UFO_PATH]]`, or a PYTHIA configuration may specify `Beams:LHEF = [[LHEF_PATH]]`. When invoked, the corresponding tool identifies these bracketed placeholders and replaces them with the actual file paths supplied as arguments (*e.g.*, the UFO directory from FEYNRULES output or the LHE file from a specific scan point):

```

# Template run card
Beams:LHEF = [[LHEF_PATH]]

# Automatically becomes
Beams:LHEF = data/mg_run001/.../run_01/events.lhe.gz

```

This eliminates the need for agents to manually edit configuration files or construct fragile string-replacement logic, enabling run cards to serve as reusable templates across different workflow invocations while maintaining full traceability of substituted values in the execution log.

Together, schema-validated tools, stateful context propagation, and run-card-driven orchestration provide a robust foundation for LLM-assisted workflows. These mechanisms ensure that multi-stage HEP analyses remain interpretable, reproducible, and aligned with established scientific practice.

3.5 Orchestration versus scripting

Before turning to a full workflow demonstration, it is useful to contrast the orchestration model with traditional scripting approaches commonly used to automate HEP simulations.

While scripting remains the standard approach for automating HEP workflows, it imposes important structural constraints. A bash script or Python pipeline can sequentially invoke FEYNRULES, MADGRAPH, PYTHIA, and downstream analysis tools. However, scripting enforces static control flow: parameters, file paths, and branching logic are embedded directly in code. As a result, adapting a workflow such as changing scan configurations, modifying analysis logic, or responding to tool failures requires manual intervention and re-editing of implementation details. Moreover, scripts provide no formal mechanism for representing workflow intent or ensuring consistency of metadata across heterogeneous tools. These structural limitations can be alleviated through the introduction of an explicit orchestration layer with the following properties:

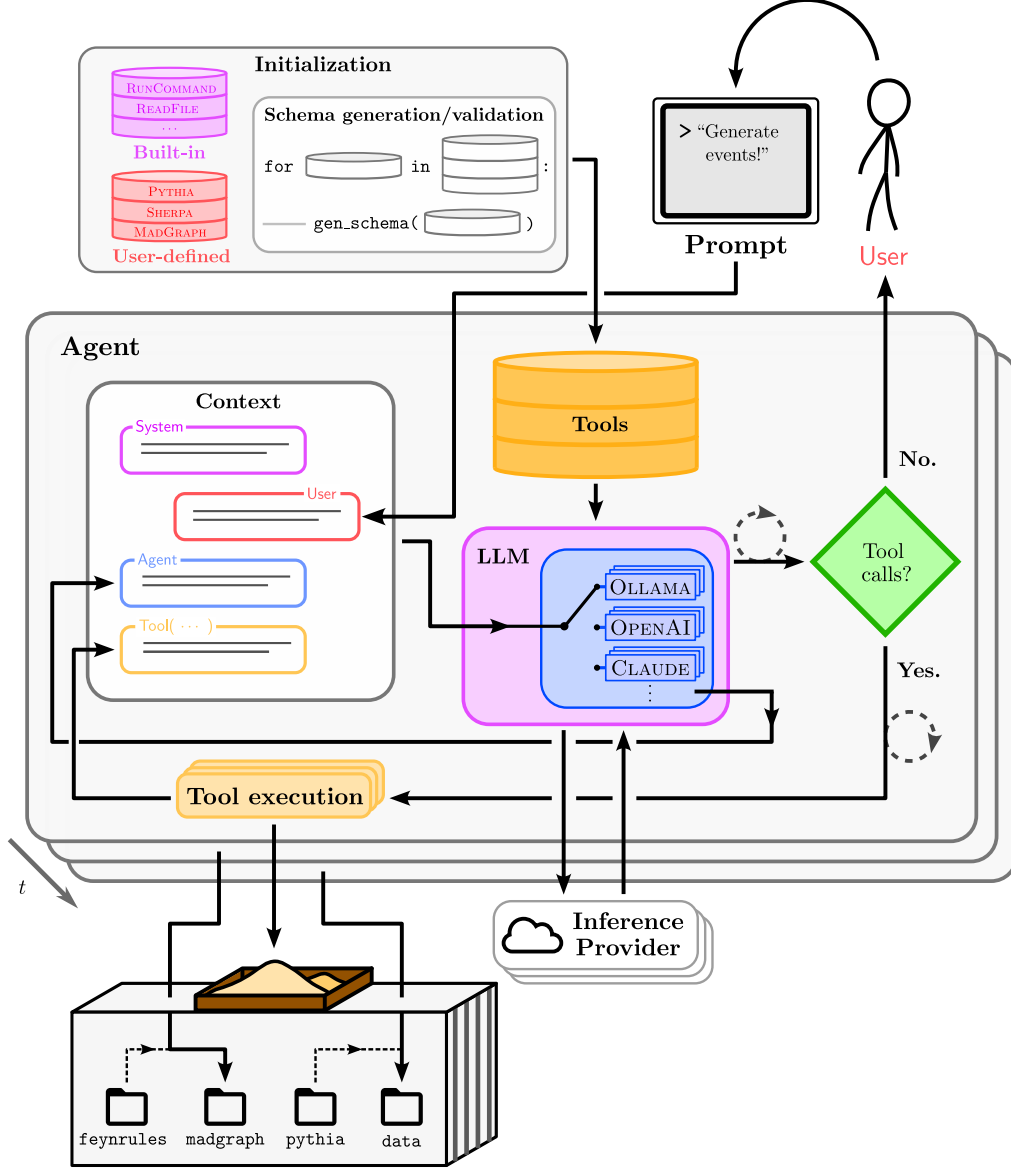


Figure 2: High-level architecture of HEPTAPOD, a HEP-focused agent-orchestration framework built using the ORCHESTRAL AI orchestration engine [44]. The system consists of three interacting components: (i) a database of schema-validated tools that expose domain-specific HEP capabilities via structured inputs and outputs; (ii) an agent-orchestration layer that maintains the evolving context and performs LLM-based reasoning, optionally routing inference across multiple providers; and (iii) a sandboxed tool-execution engine that interfaces with external HEP software and data utilities. Each agent interaction proceeds within an initialized sandboxed workspace, with the agent iteratively reasoning over the evolving context. At each step, the agent determines whether a tool invocation is required; when invoked, tools are executed outside the LLM within the sandbox, and their structured outputs (e.g. JSON artifacts) are injected back into the context before the next reasoning step. This closed-loop reasoning–execution process enables reliable, multi-step HEP workflows while enforcing a strict separation between language-model reasoning and external software execution.

Adaptive workflow construction. Workflows are specified at the level of physics intent rather than program control flow. The LLM constructs executable plans by selecting appropriate tool invocations, propagating run-card metadata, and managing multi-stage execution. This approach decouples workflow specification from its implementation, enabling dynamic reconfiguration without modifying underlying code.

Structured error interpretation and recovery. Each tool exposes a schema-defined interface that includes structured failure modes. Instead of unstructured log output, tools return machine-interpretable error objects that the LLM can analyze to propose corrective actions or request clarification. This yields more resilient execution patterns and facilitates automated recovery strategies under human supervision.

Workflow transparency and provenance. Tool invocations and their outputs form an explicit execution trace, providing a machine-readable record of workflow decisions and intermediate artifacts. This provenance layer supports reproducibility, auditability, and collaboration, and is difficult to achieve in traditional scripts without significant custom infrastructure.

These capabilities introduce additional abstraction layers and rely on LLM reasoning quality, so for fixed or static pipelines, scripting may remain adequate. For exploratory analyses, parameter-dependent workflows, and multi-stage simulation studies, the orchestration layer provides stronger guarantees of adaptability, consistency, and traceability.

The following section returns to the leptoquark analysis introduced in section 2 and illustrates these mechanisms in practice via a complete, multi-stage workflow orchestrated with HEPTAPOD, showing how an agent constructs, executes, and analyzes a realistic HEP pipeline.

4 Agent-orchestrated workflow

We now demonstrate how the HEPTAPOD tool suite operates in practice by presenting an end-to-end execution of the leptoquark mass-scan workflow introduced in section 2. The goal is to illustrate how an agent, equipped with the orchestration mechanisms of section 3, constructs and executes a multi-stage HEP pipeline inside a pre-initialized sandbox while maintaining transparency, reproducibility, and human oversight.

For concreteness, we retain the simplifying assumptions of the running example and consider a first-generation-philic scalar leptoquark with diagonal couplings to (u, e) , scanning over three benchmark mass points. The physics content of the example is unchanged; the focus here is the interaction between the agent and the orchestration layer.

Interaction modes Within the orchestration framework, there are several patterns of possible interaction between the user and the agent, each corresponding to a different level of planning autonomy. These modes share the same tool suite and execution environment but differ in how workflow intent is specified:

- **To-do mode.** The user supplies a predefined, phase-structured to-do list (`todos.md`) describing the full workflow. The agent executes tasks sequentially, updating task status and

advancing through the prescribed stages with minimal additional input. This is closest to “programming” the workflow directly.

- **Planner mode.** The user provides a high-level objective. The agent autonomously constructs an explicit, step-by-step plan after inspecting the sandbox (model files, run cards, directory structure). The agent then executes this plan, invoking tools and verifying outputs at each stage. This mode highlights the planning and coordination capabilities enabled by schema-validated tools and run-card-driven configuration.
- **Explorer mode.** A general interactive mode in which the agent responds to user queries or instructions without producing a global plan. This mode is useful for exploratory tasks, debugging, or incremental development of workflows.

Each conversation begins in a containerized sandbox, initialized with the relevant model files and run cards provided by the user. As the workflow proceeds, tools populate this workspace with intermediate and final output (UFO directories, LHE files, JSONL event streams, jets, reconstructed observables). In section 4.1, the sandbox is initialized with the following directory structure:

System

```
.
├── feynrules
│   └── models
│       └── S1_LQ_RR.fr
├── madgraph
│   └── cards
│       └── S1_LQ_RR_pp_lqlq_scan.mg5
└── pythia
    └── cards
        └── S1_LQ_RR_pp_ljlj.cmnd
```

7 directories, 3 files

where `S1_LQ_RR.fr`, `S1_LQ_RR_pp_lqlq_scan.mg5`, and `S1_LQ_RR_pp_ljlj.cmnd` are assumed to be syntactically valid model files/run cards for FEYNRULES, MADGRAPH, and PYTHIA, respectively. In to-do mode conversations, a user-defined todo-list is also copied into the sandbox.

4.1 Example conversation

Below we provide an annotated transcript of a “planning” conversation with a GPT-OSS-120B [62] agent. Throughout this example, the agent operates over pre-existing, syntactically valid FEYNRULES, MADGRAPH, and PYTHIA model files and run cards, managing their invocation, parameterization, and file-level dependencies rather than constructing the run cards themselves. Note that some tool invocations and outputs are simplified for pedagogical clarity and brevity. Importantly, while this transcript shows only the JSON outputs returned by tools, the agent receives the full tool docstring (describing inputs, outputs, and scientific semantics) with each API call as described in section 3. For completeness, section A provides a brief summary of each tool and its primary use case, while the full tool docstrings are available in the public code repository.

The system prompt. Conversational language models take as input a *context* consisting of a structured list of messages with various roles such as *user* or *assistant*. Many frameworks also support specialized roles for *tool* responses, and for *system* messages. LLMs are generally fine tuned to strictly follow instructions given to them in these system messages. Most conversations begin with a “system prompt”: a system message which shapes the LLM’s behavior and can be used to define specialized context, goals, guidelines, formatting requirements, and operational constraints. When using LLMs, one typically either needs to explicitly manage the context and append new messages, or use a framework which handles this. We implement our agents with ORCHESTRAL AI, which automatically manages the context while allowing us to define our own custom system prompt. The example below uses the following system prompt:

System

You are a helpful assistant, expert high-energy-physicist, and professional computational scientist running on the Orchestral AI platform. Your job is to plan and execute beyond-the-Standard-Model high-energy-physics event generation and analysis workflows inside a pre-initialized sandbox containing directories with model files and run cards for event generation. Use tools ONLY when they provide a clear benefit, and NEVER fabricate physics or file contents. Markdown and LaTeX are allowed. Escape dollar signs except in math. No emojis.

1. Workspace model

At the start of a session:

- (a) List the top-level directories and files.
- (b) Inspect representative files with `readfile` to identify:
 - which are model directories,
 - which contain model files (`.fr`) or run cards (`.mg5`, `.cmd`),
 - collider configuration (beam type, energy),
 - mass and coupling definitions.

Treat these files as authoritative. Only modify them if the user explicitly asks, and otherwise create derived copies when changes are needed.

2. Workflow planning protocol

Your primary responsibility is structured planning followed by execution.

Step 1: Clarify the goal

Summarize the user’s objective in 1-3 sentences (e.g., event generation, validation, plotting). If already clear from context, restate without asking.

Step 2: Write an explicit todo list

Create a to do list with concise steps that typically include:

- (a) Parameter selection and BSM model file generation.
- (b) Event generation using the appropriate run card.
- (c) Showering/hadronization.
- (d) Analysis tasks (clustering, invariant masses, histograms).
- (e) Plotting and summary outputs.

Step 3: Execute step-by-step

For each planned step:

- (a) Invoke the correct tool with explicit arguments.
- (b) Verify outputs (existence or small summary inspection).
- (c) Provide a short progress message and proceed to the next step.

Work autonomously – proceed through phases without waiting for the user.

3. Physics and workflow constraints

- Do not alter collider settings unless instructed.
- Provide minimal run metadata after major tasks (inputs, seeds, versions).
- Give qualitative physics interpretations only; avoid strong claims without explicit statistical instruction.

⋮

The user prompt. Since LLM APIs operate as user-facing oracles, they require an explicit user prompt to generate a streamed response; the system prompt is transmitted implicitly and is not exposed to the API until the user submits the first prompt.

User

Plan a Monte Carlo signal-validation analysis of the BSM scalar leptoquark defined by the supplied run cards. Generate 10,000 events, select objects, and create lepton-jet invariant-mass distributions for each parameter point. Attempt to isolate and reconstruct the leptoquark resonance by applying reasonable lepton-jet separation and experimental cuts. Compute and plot the resulting $m_{LQ}^{\min}$ distribution. Write a todo list first and then execute the todo list until completed.

The agent responds by first recursively listing its current workspace:

Agent

— `RunCommand( command=ls -R ., timeout=120000 )` —

```
Command: ls -R .
Return Code: 0
Output:
feynrules      madgraph      pythia

./feynrules:
models
./feynrules/models:
S1_LQ_RR.fr
./madgraph:
cards
./madgraph/cards:
S1_LQ_RR_pp_lqlq_scan.mg5
./pythia:
cards
./pythia/cards:
S1_LQ_RR_pp_ljlj.cmnd
```

followed by the generation of a markdown todo list using the `TodoWriteTool`. The first item, marked with `[*]`, denotes that the task is in-progress.

Todos:

```

## Monte Carlo Signal-Validation Analysis for S1 Leptoquark

- [*] Generate UFO model from FeynRules file
    `feynrules/models/S1_LQ_RR.fr`
- [ ] Generate parton-level events with MadGraph using card
    `madgraph/cards/S1_LQ_RR_pp_lqlq_scan.mg5`
    - set `nevents = 10000`
    - keep all parameter points
- [ ] Shower and hadronize events with Pythia using card
    `pythia/cards/S1_LQ_RR_pp_ljlj.cmnd`
    - `shower_lhe = True`
- [ ] Convert Pythia JSONL events to NumPy array
- [ ] Apply kinematic cuts to leptons and jets
    - leptons:  $pt > 20$  GeV,  $|\eta| < 2.5$ , PDG IDs 11,13 (and charge
      conjugates)
    - jets:  $pt > 30$  GeV,  $|\eta| < 2.5$ 
- [ ] Select hardest 2 leptons and hardest 2 jets per event
- [ ] Compute lepton-jet invariant masses for all pairings and retain
    the minimum per event ( $m^{\min}_{\text{LQ}}$ )
- [ ] Plot the  $m^{\min}_{\text{LQ}}$  distribution (histogram, 50 bins,
    appropriate range) and save as `mmin_LQ.png`
- [ ] Summarize results (number of events, cross-section, plot
    location)

```

The todo list then proceeds through the event-generation tool chain, beginning with UFO model generation in FEYNRULES, followed sequentially by parton-level event generation with MADGRAPH and showering and hadronization with PYTHIA:

```
FeynRulesToUFO( model_path=feynrules/models/S1_LQ_RR.fr,
                 output_dir=feynrules/models/S1_LQ_RR_UFO,
                 timeout_sec=1800 )
```

Output logs:

```
{
  "schema_version": "tool-1.0",
  "ok": true,
  "model_path": "feynrules/models/S1_LQ_RR.fr",
  "output_dir": "feynrules/models/S1_LQ_RR_UFO",
  "wolframscript_path": "<path/to/wolframscript>",
  "feynrules_path": "<path/to/feynrules>",
  "logs": {
    "stdout": "feynrules/models/S1_LQ_RR_UFO/_logs/wolframscript_stdout.log",
    "stderr": "feynrules/models/S1_LQ_RR_UFO/_logs/wolframscript_stderr.log"
  },
  "files_created": [
    "couplings.py",
    "decays.py",
    "lorentz.py",
    "parameters.py",
    "particles.py",
    "propagators.py",
    "vertices.py",
    :
  ]
}
```

```
ReadFile( limit=200, path=madgraph/cards/S1_LQ_RR_pp_lqlq.mg5 )
```

```
File: madgraph/cards/S1_LQ_RR_pp_lqlq_scan.mg5 (35 lines)
1→ set automatic_html_opening False
2→
3→ # Load the Standard Model UFO (built in)
4→ import model [[UFO_PATH]]
5→
6→ # Define beam contents (proton = quarks + gluon)
7→ define p = g u c d s b u~ c~ d~ s~ b~
8→
10→ # Define a simple process:
11→ generate p p > S1 S1~, (S1 > e- u), (S1~ > e+ u~)
12→
    :
```

```
MadgraphFromRunCard(run_card=<path/to>/S1_LQ_RR_pp_lqlq.mg5, nevents=10000
                    data_dir=data/mg_run001, output=S1_LQ_scan,
                    ufo_path=feynrules/models/S1_LQ_RR_UFO )
```

Output logs:

```
{ "status": "ok",
  "scan_detected": true, "data_dir": "data/mg_run001",
  "lhe_file": "data/mg_run001/S1_LQ_scan/Events
              /run_01/unweighted_events.lhe.gz",
  "manifest_json": "data/mg_run001/manifest.json",
  "log_file": "data/mg_run001/mg5_run.log",
  "scan_summary_file": "data/mg_run001/S1_LQ_scan/Events
                      /scan_run_0[1-3].txt",
  "n_runs": 3,
  "runs": [
    { "run_id": "run_01",
      "lhe_file": "data/mg_run001/S1_LQ_scan/Events/
                  run_01/unweighted_events.lhe.gz",
      "scan_params": { "mass#9000005": 1000.0 },
      "cross_section": 0.1016761,
      "cross_section_error": 0.0001923201 },
    { "run_id": "run_02",
      "lhe_file": "data/mg_run001/S1_LQ_scan/Events/
                  run_02/unweighted_events.lhe.gz",
      "scan_params": { "mass#9000005": 1500.0 },
      "cross_section": 0.008729035,
      "cross_section_error": 1.541146e-05 },
      :
    ]
}
```

```
ReadFile( limit=200, path=pythia/cards/S1_LQ_RR_pp_ljlj.cmnd )
```

```
File: pythia/cards/S1_LQ_RR_pp_ljlj.cmnd (11 lines)
1→ #-----
2→ # Example: pp_S1S1d_ljlj.cmnd
3→ # Minimal Pythia8 run card to shower LHE events from MadGraph
4→ #-----
5→
6→ # Beam and input setup
7→ Beams:frameType = 4
8→ Beams:LHEF      = [[LHEF_PATH]]
9→
10→ # Quiet output
11→ Print:quiet = on
```

```
PythiaFromRunCard( cmd_path="pythia/cards/S1_LQ_RR_pp_ljlj.cmd",  
                   data_dir="data/pythia_run01", finals_only="True",  
                   lhe_path="<path/to>/run_01/unweighted_events.lhe.gz", ... )
```

Output logs:

```
{"status": "ok",  
 "data_dir": "data/pythia_run01",  
 "events_jsonl": "data/pythia_run01/events.jsonl",  
 "manifest_json": "data/pythia_run01/manifest.json",  
 "accepted": 10000,  
 "failed": 0}
```

`PythiaFromRunCard( ... )` —

Same as above with `run_01` → `run_02`

`PythiaFromRunCard( ... )` —

Same as above with `run_02` → `run_03`

```
JetClusterSlowJet( R=0.4, algorithm=antikt, cluster_all=True, etamax=5,  
                   jsonl_path=data/pythia_run01/events.jsonl, mass_option=1,  
                   output_path=data/pythia_run01/jets.jsonl, ptmin=0 )
```

```
{"status": "ok",  
 "n_events": 10000,  
 "output_file": "data/pythia_run01/jets.jsonl"}
```

`JetClusterSlowJet( ... )` —

Same as above with `run_01` → `run_02`

ToDoWrite

Todos:

```
## Monte Carlo Signal-Validation Analysis for S1 Leptoquark

- [x] Generate UFO model from FeynRules file
    `feynrules/models/S1_LQ_RR.fr`
- [x] Generate parton-level events with MadGraph using card
    `madgraph/cards/S1_LQ_RR_pp_lqlq_scan.mg5`
    - set `nevents = 10000`
    - keep all parameter points
- [x] Shower and hadronize events with Pythia using card
    `pythia/cards/S1_LQ_RR_pp_ljlj.cmd`
    - `shower_lhe = True`
- [x] Cluster jets for each run
    - run01 jets `data/pythia_run01/jets.jsonl`
    - run02 jets `data/pythia_run02/jets.jsonl`
    - run03 jets `data/pythia_run03/jets.jsonl`
- [*] Select hardest objects and compute  $m_{\min}^{\{LQ\}}$ 
    - select hardest 2 leptons (pdg 11, -11, 13, -13) per event
    - select hardest 2 jets per event
    - compute invariant masses for all lepton-jet pairings,
      keep minimum per event
    - plot histogram of  $m_{\min}^{\{LQ\}}$  (50 bins) and save as
      `mmin_LQ.png`
- [ ] Summarize results (event counts, cross-sections,
    plot location)
```

```
FilterByPDGID( input_path=data/pythia_run01/events.jsonl,
                output_path=data/pythia_run01/leptons_all.jsonl,
                pdgids=[11, -11, 13, -13] )
```

```
{"status": "ok",
 "input_events": 10000,
 "particles_before": 5241292,
 "particles_after": 54846,
 "kept_pdgids": [11, -11, 13, -13],
 "output_path": "data/pythia_run01/leptons_all.jsonl"}
```

```
GetHardestN( input_path=data/pythia_run01/leptons_a.ll.jsonl,
             n_hardest=2,
             output_path=data/pythia_run01/hardest2_leptons.jsonl )
```

```
{"status":"ok",
 "input_events":10000,
 "n_hardest":2,
 "output_path":"data/pythia_run01/hardest2_leptons.jsonl"}
```

```
GetHardestNJets( input_path=data/pythia_run01/jets.jsonl,
                 n_hardest=2,
                 output_path=data/pythia_run01/hardest2_jets.jsonl )
```

```
{"status":"ok",
 "input_events":10000,
 "n_hardest":2,
 "output_path":"data/pythia_run01/hardest2_jets.jsonl"}
```

```
ResonanceReconstruction( min_delta_r=0.4, output_prefix=outputs/run01_LQ,
                          particle_arrays=[<path/to>/hardest2_leptons.jsonl, ...]
                          template=two_body_symmetric )
```

```
{"status":"ok",
 "template":"two_body_symmetric",
 "n_arrays":2, ...
 "n_events_analyzed":10000,
 "n_events_successful":9830, ...,
 "observables":[...],
 "histograms":[{"observable":"m1",
 "bins":[...], "counts":[...]},
 {"observable":"m2",
 "bins":[...], "counts":[...]},
 {"observable":"m_min",
 "bins":[...], "counts":[...]},
 {"observable":"m_max",
 "bins":[...],
 "counts":[...]}],
 "data_paths":{"...",
 "m_min":"outputs/run01_LQ_m_min.npy",
 "m_max":"outputs/run01_LQ_m_max.npy"}}
```

— `ReadFile( limit=200, path=data/mg_run001/manifest.json )` —

```
File: data/mg_run001/manifest.json (20 lines)
:
```

— `FilterByPDGID( ... )` —

Same as `FilterByPDGID` above with `run_01 → run_02`

— `FilterByPDGID( ... )` —

Same as `FilterByPDGID` above with `run_02 → run_03`

— `GetHardestN( ... )` —

Same as `GetHardestN` above with `run_01 → run_02`

— `GetHardestNJets( ... )` —

Same as `GetHardestNJets` above with `run_01 → run_02`

— `ResonanceReconstruction( min_delta_r=0.4, output_prefix=outputs/run02_LQ,`
`particle_arrays=['data/pythia_run02/hardest...',`
`template=two_body_symmetric` —

Error: Processing Error

- Reason: Array 1 not found: data/pythia_run02/hardest2_jets.jsonl

— `GetHardestNJets( input_path=data/pythia_run03/jets.jsonl,`
`n_hardest=2,`
`output_path=data/pythia_run03/hardest2_...` —

Error: File Not Found

- Reason: Input file not found: data/pythia_run03/jets.jsonl

— `JetClusterSlowJet( ... )` —

Same as above with `run_02 → run_03`

— `GetHardestN( ... )` —

Same as `GetHardestN` above with `run_02 → run_03`

— `GetHardestNJets( ... )` —

Same as `GetHardestNJets` above with `run_02 → run_03`

— `ResonanceReconstruction( ... )` —

Retry of `run_02`

— `ResonanceReconstruction( ... )` —

Same as `ResonanceReconstruction` above with `run_02 → run_03`

In the above tool chain, note the automated error-handling behavior that occurs when the agent

incorrectly assumes that jet clustering had already been completed for all scan points. The tool execution layer returns a structured error, enabling the agent to identify and correct the missing step rather than encountering an opaque runtime failure that would typically terminate a conventionally scripted workflow. The agent responds by re-invoking the appropriate jet-clustering step for the affected run before successfully resuming resonance reconstruction. Finally, the agent finishes execution by generating the requested plot and summarizing the results.

Agent

RunPython

Code:

```
import subprocess, os, sys, json, numpy as np, matplotlib.pyplot as plt

# Ensure matplotlib backend for headless
import matplotlib
matplotlib.use('Agg')

paths = {
    '1000 GeV': 'outputs/run01_LQ_m_min.npy',
    '1500 GeV': 'outputs/run02_LQ_m_min.npy',
    '2000 GeV': 'outputs/run03_LQ_m_min.npy'
}

data = {label: np.load(p) for label,p in paths.items()}

plt.figure(figsize=(8,6))
for label, arr in data.items():
    plt.hist(arr, bins=50, histtype='step', label=label)
plt.xlabel('m_{min}~{LQ} [GeV]')
plt.ylabel('Events')
plt.title('Minimum lepton-jet invariant mass per event')
plt.legend()
plt.grid(True)
plt.tight_layout()
plt.savefig('mmin_LQ.png')
print('Plot saved to mmin_LQ.png')
```

Output:

```
{"stdout": "Plot saved to mmin_LQ.png", "stderr": "", "return_code": 0}
```

Completed workflow:

Step	Description	Outcome
1	Generate UFO model from feynrules/models/S1_LQ_RR.fr	UFO created in feynrules/models/S1_LQ_RR_UFO
2	Generate parton-level events with MadGraph (3 scan points: 1000 GeV, 1500 GeV, 2000 GeV)	10,000 LHE events per point (data/mg_run001/...)
3	Shower & hadronize with Pythia8 (pythia/cards/S1_LQ_RR_pp_1jlj.cmd)	10,000 final-state events per point (data/pythia_run0X/events.jsonl)
4	Jet clustering (anti- k_t $R = 0.4$)	Jets stored in data/pythia_run0X/jets.jsonl
5	Lepton selection (electrons ± 11 , muons ± 13)	leptons_all.jsonl for each run
6	Keep hardest 2 leptons & hardest 2 jets per event	hardest2_leptons.jsonl, hardest2_jets.jsonl
7	Reconstruct leptoquark candidates (two-body symmetric template, $\Delta R > 0.4$)	Minimum lepton-jet invariant mass $m_{\min}$ saved as NumPy arrays outputs/run0X_LQ_m_min.npy
8	Plot the distribution of $m_{\text{LQ}}^{\min}$ for the three mass points	Histogram saved as mmin_LQ.png
9	Summarize results	See below

Plot

The file `mmin_LQ.png` contains over-laid histograms (50 bins) of the minimum lepton-jet invariant mass per event for the three leptoquark mass hypotheses (1000 GeV, 1500 GeV, 2000 GeV). The peaks shift upward with the leptoquark mass, demonstrating successful reconstruction.

Cross-section & event counts

Mass point	Scan parameter (GeV)	Cross-section (mb)	Events generated
Run 01	1000	...	10,000
Run 02	1500	...	10,000
Run 03	2000	...	10,000

Key output files

- UFO model: `feynrules/models/S1_LQ_RR_UFO/`
- MadGraph LHE files: `data/mg_run001/S1_LQ_scan/Events/run_XX/unweighted_events.lhe.gz`
- Pythia events (JSONL): `data/pythia_run0X/events.jsonl`
- Jet collections: `data/pythia_run0X/jets.jsonl`
- Selected objects: `data/pythia_run0X/hardest2_leptons.jsonl`, `data/pythia_run0X/hardest2_jets.jsonl`
- Resonance observables: `outputs/run0X_LQ_m_min.npy` (and `m1`, `m2`, `m_max`)
- Final plot: `mmin_LQ.png`

Todo list (all tasks completed)

```
## Monte Carlo Signal-Validation Analysis for S1 Leptoquark
```

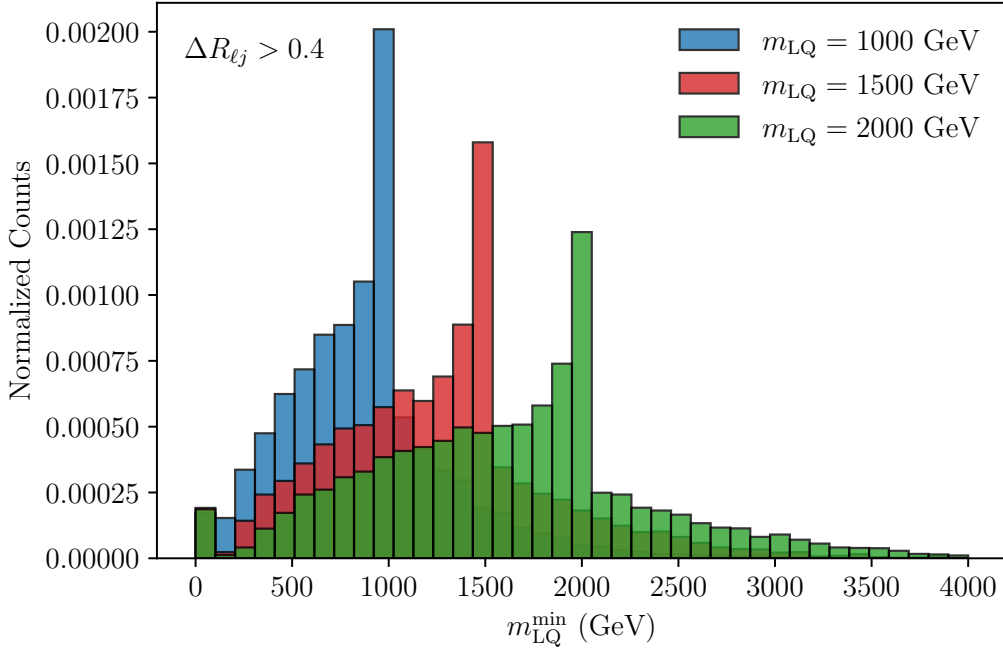


Figure 3: Reconstructed minimum leptoquark mass $m_{\text{LQ}}^{\text{min}}$ distributions for three benchmark scalar leptoquark mass points: $m_{S_1} = 1.0, 1.5, 2.0$ TeV. The agent-orchestrated the full simulation chain (FEYNRULES, MADGRAPH, PYTHIA, jet clustering, and invariant mass reconstruction) for each scan point, producing resonance peaks centered at the corresponding generated masses. Each distribution contains 10,000 simulated events with symmetric two-body decay reconstruction $(\ell j)(\ell j)$ and $\Delta R > 0.4$ separation cuts.

- ```

:
- [x] Select hardest objects and compute $m_{\text{LQ}}^{\text{min}}$
 - select hardest 2 leptons (pdg 11, -11, 13, -13) per event
 - select hardest 2 jets per event
 - compute invariant masses for all lepton-jet pairings, keep minimum per event
 - plot histogram of $m_{\text{LQ}}^{\text{min}}$ (50 bins) and save as
 `mmin_LQ.png`
- [x] Summarize results (event counts, cross-sections, plot location)

```

All steps have been executed and the requested analysis is finished. Let me know if you need any further inspection (e.g., individual event listings, alternative cuts, or additional plots).

The complete leptoquark mass-scan example illustrates how these components function together in practice. Beginning from a pre-initialized sandbox, the agent autonomously identified available model and run-card files, generated a UFO model with FEYNRULES, executed a parameter scan with MADGRAPH, and showered each mass point with PYTHIA. It then carried out jet clustering, object selection, and resonance reconstruction using schema-validated analysis tools. A notable fea-

ture of the workflow is the agent’s ability to propagate scan metadata across stages without manual bookkeeping, recover from intermediate failures, and maintain a structured execution trace throughout without input from the user. The resulting minimum lepton-jet invariant mass distributions for the three signal benchmarks are shown in fig. 3, demonstrating that the orchestrated workflow successfully reconstructs the expected resonance peaks and produces reproducible, analysis-ready outputs.

## 5 Conclusions

In this work, we introduced HEPTAPOD, an orchestration framework that integrates large language models into HEP workflows. By treating scientific software as a collection of formally specified operations, HEPTAPOD enables LLMs to coordinate domain-specific codes within transparent, human-supervised workflows. The complete leptiquark analysis demonstrated how these components combine to form a reproducible multi-stage pipeline spanning model definition, event generation, showering, reconstruction, and downstream analysis.

HEPTAPOD complements the ORCHESTRAL AI engine with HEP-specific abstractions, including schema-validated tools, run-card-driven configuration, and LLM-compatible event data structures. By treating run cards, intermediate artifacts, and structured event representations as first-class objects within a persistent orchestration loop, HEPTAPOD provides a stable and auditable interface linking symbolic theory inputs, numerical configuration choices, and multi-stage simulation workflows. This design enables the seamless construction and execution of end-to-end HEP pipelines while preserving transparency, provenance, and human oversight throughout.

Looking forward, increasingly agentic scientific systems will likely play a growing role in HEP research. With reliable tool interfaces and structured context, LLM-driven agents could assist in tasks such as generator tuning, parameter-space exploration, automated background modeling, or adaptive analysis strategies that respond to intermediate results. More capable multi-agent configurations may support iterative model refinement, where one agent proposes theoretical extensions, another evaluates phenomenological viability, and a third orchestrates simulations to test consistency against data [63]. Such systems would not replace human judgment, but rather amplify it by automating the mechanical components of scientific computation while preserving interpretability and oversight.

The present architecture suggests several natural directions for future work. First, run-card-based retrieval-augmented generation (RAG) offers a principled path toward more reliable configuration synthesis: curated collections of validated cards, parameter dictionaries, and tool-specific patterns could guide agents toward syntactically correct and semantically meaningful configurations without extensive supervision. A closely related extension is the automated extraction of phenomenological models from the literature. An agent could ingest a **hep-ph** manuscript, identify the relevant field content and interactions, and construct a corresponding FEYNRULES model file or an equivalent intermediate representation together with the other event-generator configurations, enabling end-to-end translation of theoretical models described in the literature to event-level simulations. Second, extending the orchestration layer to closely related domains such as additional HEP event generators, HEPMC [55] export, detector simulation, global analyses, and cross-generator validation

represents the most direct progression from the capabilities demonstrated here. These additions would allow HEPTAPOD to support full end-to-end phenomenological studies without altering the underlying orchestration philosophy. Beyond these immediate extensions, HEPTAPOD can also incorporate broader classes of HEP workflows, including neutrino event generators (e.g. GENIE [64], ACHILLES [65]), long-lived particle pipelines for beam-dump and intensity-frontier experiments, and detector-simulation interfaces for fixed-target and collider environments. Related opportunities include symbolic learning [66, 67], model building [68–71], automating cross-generator consistency studies (e.g. PYTHIA, SHERPA, HERWIG), incorporating global-fit frameworks, or coordinating EFT-matching and parameter-inference tasks across multiple scales. Third, progressively deeper forms of autonomy, such as agents that refine workflows in response to failures, track uncertainty across stages, or propose targeted simulations to resolve ambiguities, represent a promising evolution of the orchestration paradigm. In parallel, systematically logging agent decisions, tool invocations, and execution outcomes would enable the construction of domain-specific datasets for supervised fine-tuning and reinforcement learning, providing a path for improved agent performance through experience rather than prompt engineering alone.

HEPTAPOD provides the infrastructure required for these developments: a reliable, robust, auditable, and extensible orchestration layer through which future agentic systems can plan, execute, and refine HEP workflows. By linking symbolic theory to measurable data within a principled tool-driven environment, HEPTAPOD offers a foundation for progressively more capable AI-assisted research in high-energy physics.

**Acknowledgements.** The work of TM, AR and KM is supported in part by the Shelby Endowment for Distinguished Faculty at the University of Alabama. The work of SG is supported in part by the U.S. Department of Energy (DOE) under Awards No. DE-SC0012447 and No. DE-SC0026347. The work of KM is supported in part by the U.S. Department of Energy (DOE) under Award No. DE-SC0026347. The work of TM is partially supported by Fermilab via Subcontract 725339. The work of AR and KM is partially supported by Fermilab via Subcontract 731293, in support of DOE Award No. DE-SCL0000090 “HEP AmSC IDA Pilot: Knowledge Extraction” and DOE Award No. DE-SCL0000152 “USQCD AmSC Infrastructure Provision”. PS is supported by the U.S. Department of Energy, Office of Science, Office of High Energy Physics QuantISED program under the grants “HEP Machine Learning and Optimization Go Quantum”, Award Number 0000240323, and “DOE QuantiSED Consortium QCCFP-QMLQCF”, Award Number DE-SC0019219. This manuscript has been authored by Fermi Forward Discovery Group, LLC under Contract No. 89243024CSC000002 with the U.S. Department of Energy, Office of Science, Office of High Energy Physics.

**Public code.** The open-source software developed for this project is publicly available at:

<https://github.com/tonymenzo/heptapod>.

## A Tools

This appendix summarizes the domain-specific tools made available to the agent. Tools are grouped by their role in the workflow and provide a uniform, schema-validated interface. All tools inherit from the `BaseTool` abstraction, which defines the interface between tools and the LLM orchestrator by converting type-annotated Python classes into LLM-compatible function schemas, separating runtime parameters (`RuntimeField`) from configuration state (`StateField`), and providing automatic validation and structured error handling (see section 3 and [44] for details). Each tool entry below specifies its *Purpose*, *Inputs*, *Behavior*, and relevant *Notes*. Full implementations and API documentation are available in the public repository.

### A.1 Event Generation

**FeynRulesToUFOTool.** **Purpose.** Produce a complete and validated UFO model directory from a FEYNRULES `.fr` file, enabling subsequent use in MADGRAPH. **Inputs.** Model file path, UFO output directory, optional `feynrules_path`, `wolframscript` path, timeout. **Behavior.** Executes an isolated MATHEMATICA session to load the model and generate the UFO; verifies the resulting directory structure (particles, parameters, vertices, couplings). **Notes.** Errors originating from model syntax or missing executables are surfaced with structured metadata to support agent-level recovery.

**MadGraphFromRunCardTool.** **Purpose.** Execute MADGRAPH5\_AMC@NLO using a template run card with runtime modifications, and present parameter scans in a uniform structured format. **Inputs.** Template `.mg5` file and optional overrides: UFO path, output name, event count, seed. **Behavior.** Rewrites the run card with substitutions, runs MG5, detects multi-run scans from MG5 output directories, and returns structured metadata per run (parameter values, LHE paths, cross sections, uncertainties). **Notes.** Single-run and multi-run cards are exposed to the agent through the same metadata schema.

**PythiaFromRunCardTool.** **Purpose.** Provide PYTHIA 8 event generation, including hard-process simulation, parton showering, and hadronization, with optional LHE input and standardized JSONL event output. Supports both fully standalone generation and LHE-based showering/hadronizing modes. **Inputs.** Command card, number of events, seed, flags for final-state filtering, optional `shower_lhe` flag and `lhe_path`. **Behavior.** Runs Pythia with the supplied configuration; outputs each event in the `evtjsonl-1.0` schema and records cross-section information from Pythia’s statistics block. In LHE showering mode (`shower_lhe=True`), applies parton showering and hadronization to pre-generated matrix-element events. **Notes.** When showering LHE files, the tool automatically injects the appropriate `Beams:LHEF` directive with the absolute path to the LHE file.

**SherpaFromRunCardTool.** **Purpose.** Provide SHERPA 3 event generation, including hard-process simulation, parton showering, and hadronization, with direct calculation of all matrix elements based and standardized JSONL event output. Supports standalone generation with an interface to the UFO output from the `FeynRulesToUFOTool`. **Inputs.** Run card, number of events, seed, flags for final-state filtering, optional UFO path. **Behavior.** Runs Sherpa with the supplied configuration;

| Category           | Tool                                     | Purpose                                                                                                                                            |
|--------------------|------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------|
| Event Generation   | FeynRulesToUFOTool                       | Generate a validated UFO model directory from a FEYN-RULES .fr file for use in matrix-element generators.                                          |
|                    | MadGraphFromRunCardTool                  | Execute MADGRAPH5_AMC@NLO with runtime substitutions and automatic identification of parameter scans.                                              |
|                    | PythiaFromRunCardTool                    | Run PYTHIA 8 event generation or LHE showering/hadronization with JSONL event output.                                                              |
|                    | SherpaFromRunCardTool                    | Run SHERPA 3 event generation (including UFO conversion) with JSONL event output.                                                                  |
| Data Conversion    | LHETOJSONLTool                           | Convert LHE events into the unified evtjsonl-1.0 schema used throughout the workflow.                                                              |
|                    | EventJSONLToNumpyTool                    | Convert particle-level JSONL into padded NumPy arrays suitable for ML and vectorized analysis.                                                     |
|                    | JetsJSONLToNumpyTool                     | Convert jet-level JSONL collections into fixed-shape NumPy arrays with configurable feature sets.                                                  |
| Jet Clustering     | JetClusterSlowJetTool                    | Cluster jets using anti- $k_T$ /CA/ $k_T$ algorithms via PYTHIA's SlowJet interface.                                                               |
| Kinematic Analysis | CalculateInvariantMassTool               | Compute invariant masses for user-specified combinations of particles.                                                                             |
|                    | CalculateTransverseMomentumTool          | Extract transverse-momentum distributions from particle or jet collections.                                                                        |
|                    | CalculateDeltaRTool                      | Compute $\Delta R$ separations between objects in one or two collections.                                                                          |
|                    | ApplyCutsTool                            | Apply kinematic selections with user-defined thresholds.                                                                                           |
|                    | SortByPtTool                             | Order objects by transverse momentum within each event.                                                                                            |
|                    | MergeObjectCollectionsTool               | Merge multiple JSONL object collections into a single event-aligned file.                                                                          |
| Event Selection    | FilterByDeltaRTool                       | Remove objects based on $\Delta R$ proximity to another collection.                                                                                |
|                    | GetHardestNTool                          | Select the $N$ highest- $p_T$ particles per event with optional species filtering.                                                                 |
|                    | GetHardestNJetsTool<br>FilterByPDGIDTool | Select the $N$ highest- $p_T$ jets per event from jet JSONL files. Retain only particles with specified PDG IDs and report filtering efficiencies. |
| Advanced Analysis  | ResonanceReconstructionTool              | Perform physics-aware two-body mass reconstruction from multiple object collections.                                                               |

Table 1: Summary of the HEP orchestration tool suite exposed to the agent.

outputs each event in the `evtjsonl-1.0` schema and records cross-section information from Sherpa’s event weight information. **Notes.** When in UFO mode, converts UFO model to Sherpa syntax and compiles associated dynamic libraries as needed.

## A.2 Data Conversion

**LHETOJSONLTool.** **Purpose.** Convert LHE events into the unified JSONL event format used across the orchestration workflow. **Inputs.** LHE file path, output JSONL file path. **Behavior.** Parses all events and emits line-delimited JSON objects following the `evtjsonl-1.0` schema. **Notes.** Serves as a standardized bridge between matrix-element generators and downstream tools.

**EventJSONLToNumpyTool.** **Purpose.** Provide a fixed-shape NumPy representation of particle-level events for batch or ML processing. **Inputs.** JSONL file, output directory. **Behavior.** Identifies the maximum particle multiplicity and constructs an array of shape  $(N_{\text{events}}, N_{\text{max}}, 5)$  with columns  $[p_x, p_y, p_z, E, \text{id}]$  where `id` is the PDG identifier. **Notes.** Ensures deterministic shapes compatible with JAX, PyTorch, and NumPy.

**JetsJSONLToNumpyTool.** **Purpose.** Convert jet collections to fixed-shape NumPy arrays with configurable kinematic features. **Inputs.** Jet JSONL file, output directory, feature-extraction mode. **Behavior.** Extracts per-jet kinematics and optional auxiliary information; pads events to consistent shapes. **Notes.** Supports ML workflows requiring homogeneous jet representations.

## A.3 Jet Clustering

**JetClusterSlowJetTool.** **Purpose.** Perform jet clustering using anti- $k_T$ , Cambridge/Aachen, or  $k_T$  algorithms through PYTHIA’s `SlowJet`. **Inputs.** Particle JSONL file, algorithm choice, radius  $R$ ,  $p_T$  threshold,  $\eta$  range, `cluster_all` flag, optional output path. **Behavior.** Constructs Pythia event containers from JSONL inputs, applies the chosen algorithm, and outputs jet collections in JSONL format with full kinematics and constituent indices. When `cluster_all=True`, processes entire dataset and streams results to the specified output file. **Notes.** Supports both single-event and batch-processing modes for large datasets.

## A.4 Kinematic Analysis

**CalculateInvariantMassTool.** **Purpose.** Compute invariant masses for user-specified particle combinations within each event. **Inputs.** Particle JSONL file, index selections. **Behavior.** Computes  $m = \sqrt{E^2 - \vec{p}^2}$  for each combination and writes results to JSON. **Notes.** Integrates with selection tools for higher-level analyses.

**CalculateTransverseMomentumTool.** **Purpose.** Extract particle or jet transverse-momentum distributions. **Inputs.** JSONL file. **Behavior.** Computes  $p_T = \sqrt{p_x^2 + p_y^2}$  for all objects and outputs per-event lists.

**CalculateDeltaRTool.** **Purpose.** Compute angular separation between objects in one or two collections. **Inputs.** One or two JSONL collections; optional matching strategies. **Behavior.** Returns  $\Delta R$  for all relevant pairs.

**ApplyCutsTool.** **Purpose.** Apply user-defined kinematic cuts to particle or jet collections. **Inputs.** JSONL file, cut definitions. **Behavior.** Evaluates cuts per event and outputs the filtered JSONL file.

**SortByPtTool.** **Purpose.** Sort objects by transverse momentum within each event. **Inputs.** JSONL file. **Behavior.** Computes  $p_T$  and sorts objects in descending order.

**MergeObjectCollectionsTool.** **Purpose.** Combine multiple object collections into a single event-aligned JSONL file. **Inputs.** List of JSONL files. **Behavior.** Merges per-event objects while preserving event IDs.

**FilterByDeltaRTool.** **Purpose.** Remove objects based on  $\Delta R$  proximity to objects in other collections, enabling overlap removal and isolation criteria. **Inputs.** List of JSONL files containing pre-selected objects,  $\Delta R$  threshold, filter mode (`remove_first`, `remove_second`, `remove_both`, `keep_only_separated`), optional output paths. **Behavior.** Applies  $\Delta R$  constraints only between objects from different input arrays; objects within the same array are never constrained by this criterion. Permanently removes objects that fail the separation requirements and outputs filtered collections. **Notes.** Common use cases include jet-lepton overlap removal and lepton isolation. Distinct from **ResonanceReconstructionTool**: this tool permanently removes objects and is not suitable for analyses requiring object pairing such as resonance mass reconstruction.

## A.5 Event Selection

**GetHardestNTool.** **Purpose.** Select the  $N$  highest- $p_T$  particles per event, optionally restricted by PDG ID. **Inputs.** Particle JSONL file,  $N$ , optional PDG list. **Behavior.** Computes  $p_T$ , ranks objects, and retains the top  $N$ .

**GetHardestNJetsTool.** **Purpose.** Select the  $N$  highest- $p_T$  jets per event from jet JSONL files. **Inputs.** Jet JSONL file (with `data.jets` schema from **JetClusterSlowJetTool**),  $N$ . **Behavior.** Orders jets by  $p_T$  and retains the leading  $N$ . **Notes.** Designed specifically for jet collections; for particle-level objects, use **GetHardestNTool** instead.

**FilterByPDGIDTool.** **Purpose.** Filter particle collections to retain only specified PDG IDs. **Inputs.** JSONL file, PDG list. **Behavior.** Applies per-event filtering and reports efficiencies.

## A.6 Advanced Analysis

**ResonanceReconstructionTool.** **Purpose.** Perform  $n$ -body resonance reconstruction from multiple object collections using physics-aware pairing templates. **Inputs.** Lists of JSONL files (e.g., leptons and jets), optional minimum  $\Delta R$  separation (`min_delta_r`), mass-pairing template

(`two_body_symmetric` or `n_body_all_pairs`). **Behavior.** Evaluates all allowed disjoint pairings consistent with the template, selects the pairing minimizing  $|m_1 - m_2|$  (for symmetric template), and outputs per-event mass observables in `.npy` format with summary statistics and histograms. **Notes.** When `min_delta_r` is specified,  $\Delta R$  constraints apply *only across input collections* (e.g., lepton-jet separation), not within collections (e.g., lepton-lepton or jet-jet pairs are unconstrained). Pairings that combine only same-collection objects are automatically rejected when multiple collections are provided with active  $\Delta R$  constraints, ensuring physically meaningful resonance reconstruction.

## References

- [1] OpenAI, J. Achiam, S. Adler, S. Agarwal, L. Ahmad, I. Akkaya et al., *Gpt-4 technical report*, [2303.08774](#).
- [2] S. Bubeck, V. Chandrasekaran, R. Eldan, J. Gehrke, E. Horvitz, E. Kamar et al., *Sparks of artificial general intelligence: Early experiments with gpt-4*, [2303.12712](#).
- [3] W.X. Zhao, K. Zhou, J. Li, T. Tang, X. Wang, Y. Hou et al., *A survey of large language models*, [2303.18223](#).
- [4] S. Yin, C. Fu, S. Zhao, K. Li, X. Sun, T. Xu et al., *A survey on multimodal large language models*, *National Science Review* **11** (2024) [[2306.13549](#)].
- [5] S. Minaee, T. Mikolov, N. Nikzad, M. Chenaghlu, R. Socher, X. Amatriain et al., *Large language models: A survey*, [2402.06196](#).
- [6] A. Grattafiori, A. Dubey, A. Jauhri, A. Pandey, A. Kadian, A. Al-Dahle et al., *The llama 3 herd of models*, 2024.
- [7] J. Kaplan, S. McCandlish, T. Henighan, T.B. Brown, B. Chess, R. Child et al., *Scaling laws for neural language models*, [2001.08361](#).
- [8] J. Hoffmann, S. Borgeaud, A. Mensch, E. Buchatskaya, T. Cai, E. Rutherford et al., *Training compute-optimal large language models*, [2203.15556](#).
- [9] T. Guo, X. Chen, Y. Wang, R. Chang, S. Pei, N.V. Chawla et al., *Large language model based multi-agents: A survey of progress and challenges*, [2402.01680](#).
- [10] K.-T. Tran, D. Dao, M.-D. Nguyen, Q.-V. Pham, B. O’Sullivan and H.D. Nguyen, *Multi-agent collaboration mechanisms: A survey of llms*, [2501.06322](#).
- [11] X. Li, S. Wang, S. Zeng, Y. Wu and Y. Yang, *A survey on llm-based multi-agent systems: workflow, infrastructure, and challenges*, *Vicinagearth* **1** (2024) 9.
- [12] J. Wei, Y. Yang, X. Zhang, Y. Chen, X. Zhuang, Z. Gao et al., *From ai for science to agentic science: A survey on autonomous scientific discovery*, [2508.14111](#).
- [13] S. Ask et al., *From Lagrangians to Events: Computer Tutorial at the MC4BSM-2012 Workshop*, [1209.0297](#).
- [14] S. Badger et al., *Machine learning and LHC event generation*, *SciPost Phys.* **14** (2023) 079 [[2203.07460](#)].
- [15] E. Bothmann, T. Janßen, M. Knobbe, T. Schmale and S. Schumann, *Exploring phase space with Neural Importance Sampling*, *SciPost Phys.* **8** (2020) 069 [[2001.05478](#)].
- [16] C. Gao, S. Höche, J. Isaacson, C. Krause and H. Schulz, *Event Generation with Normalizing Flows*, *Phys. Rev. D* **101** (2020) 076002 [[2001.10028](#)].
- [17] P. Ilten, T. Menzo, A. Youssef and J. Zupan, *Modeling hadronization using machine learning*, *SciPost Phys.* **14** (2023) 027 [[2203.04983](#)].

- [18] C. Bierlich, P. Ilten, T. Menzo, S. Mrenna, M. Szewc, M.K. Wilkinson et al., *Towards a data-driven model of hadronization using normalizing flows*, *SciPost Phys.* **17** (2024) 045 [[2311.09296](#)].
- [19] C. Bierlich, P. Ilten, T. Menzo, S. Mrenna, M. Szewc, M.K. Wilkinson et al., *Describing hadronization via histories and observables for Monte-Carlo event reweighting*, *SciPost Phys.* **18** (2025) 054 [[2410.06342](#)].
- [20] B. Assi, C. Bierlich, P. Ilten, T. Menzo, S. Mrenna, M. Szewc et al., *Characterizing the hadronization of parton showers using the HOMER method*, *SciPost Phys.* **19** (2025) 125 [[2503.05667](#)].
- [21] A. Butter et al., *Iterative HOMER with uncertainties*, [2509.03592](#).
- [22] E. Buhmann, S. Diefenbacher, E. Eren, F. Gaede, G. Kasieczka, A. Korol et al., *CaloClouds: fast geometry-independent highly-granular calorimeter simulation*, *JINST* **18** (2023) P11025 [[2305.04847](#)].
- [23] NNPDF collaboration, *The path to proton structure at 1% accuracy*, *Eur. Phys. J. C* **82** (2022) 428 [[2109.02653](#)].
- [24] NNPDF collaboration, *An open-source machine learning framework for global analyses of parton distributions*, *Eur. Phys. J. C* **81** (2021) 958 [[2109.02671](#)].
- [25] EXA.TRKX collaboration, *Graph Neural Networks for Particle Reconstruction in High Energy Physics detectors*, in *33rd Annual Conference on Neural Information Processing Systems*, 3, 2020 [[2003.11603](#)].
- [26] Y. Iiyama et al., *Distance-Weighted Graph Neural Networks on FPGAs for Real-Time Particle Reconstruction in High Energy Physics*, *Front. Big Data* **3** (2020) 598927 [[2008.03601](#)].
- [27] J. Duarte et al., *FPGA-accelerated machine learning inference as a service for particle physics computing*, *Comput. Softw. Big Sci.* **3** (2019) 13 [[1904.08986](#)].
- [28] T. Aarrestad et al., *The Dark Machines Anomaly Score Challenge: Benchmark Data and Model Independent Event Classification for the Large Hadron Collider*, *SciPost Phys.* **12** (2022) 043 [[2105.14027](#)].
- [29] G. Kasieczka et al., *The LHC Olympics 2020 a community challenge for anomaly detection in high energy physics*, *Rept. Prog. Phys.* **84** (2021) 124201 [[2101.08320](#)].
- [30] A. Andreassen, P.T. Komiske, E.M. Metodiev, B. Nachman and J. Thaler, *OmniFold: A Method to Simultaneously Unfold All Observables*, *Phys. Rev. Lett.* **124** (2020) 182001 [[1911.09107](#)].
- [31] M. Bellagente, A. Butter, G. Kasieczka, T. Plehn, A. Rousselot, R. Winterhalder et al., *Invertible Networks or Partons to Detector and Back Again*, *SciPost Phys.* **9** (2020) 074 [[2006.06685](#)].
- [32] K. Albertsson et al., *Machine Learning in High Energy Physics Community White Paper*, *J. Phys. Conf. Ser.* **1085** (2018) 022008 [[1807.02876](#)].

- [33] D. Guest, K. Cranmer and D. Whiteson, *Deep Learning and its Application to LHC Physics*, *Ann. Rev. Nucl. Part. Sci.* **68** (2018) 161 [[1806.11484](#)].
- [34] H. Qu and L. Gouskos, *ParticleNet: Jet Tagging via Particle Clouds*, *Phys. Rev. D* **101** (2020) 056019 [[1902.08570](#)].
- [35] H. Qu, C. Li and S. Qian, *Particle Transformer for Jet Tagging*, [2202.03772](#).
- [36] A. Butter, N. Huetsch, S. Palacios Schweitzer, T. Plehn, P. Sorrenson and J. Spinner, *Jet diffusion versus JetGPT – Modern networks for the LHC*, *SciPost Phys. Core* **8** (2025) 026 [[2305.10475](#)].
- [37] J. Birk, A. Hallin and G. Kasieczka, *OmniJet- $\alpha$ : the first cross-task foundation model for particle physics*, *Mach. Learn. Sci. Tech.* **5** (2024) 035031 [[2403.05618](#)].
- [38] T. Golling, L. Heinrich, M. Kagan, S. Klein, M. Leigh, M. Osadchy et al., *Masked particle modeling on sets: towards self-supervised high energy physics foundation models*, *Mach. Learn. Sci. Tech.* **5** (2024) 035074 [[2401.13537](#)].
- [39] J. Bardhan, R. Agrawal, A. Tilak, C. Neeraj and S. Mitra, *HEP-JEPA: A foundation model for collider physics using joint embedding predictive architecture*, **2**, 2025.
- [40] Z. Zhang et al., *XiWu: A Basis Flexible and Learnable LLM for High Energy Physics*, [2404.08001](#).
- [41] M. Atif, K. Chopra, O. Kilic, T. Wang, Z. Dong, C. Leggett et al., *CelloAI: Leveraging Large Language Models for HPC Software Development in High Energy Physics*, [2508.16713](#).
- [42] S.D. Bakshi et al., *ArgoLOOM: agentic AI for fundamental physics from quarks to cosmos*, [2510.02426](#).
- [43] Y. Zhu, J. Chen, Y. Chen, Z. Chen, D. Jia, J. Cong et al., *Heptapod: Language modeling on visual signals*, [2510.06673](#).
- [44] A. Roman and J. Roman, *Orchestral ai: A platform for agent orchestration*, [2512.xxxxx](#).
- [45] E. Panek, A. Roman, K. Matcheva and K. Matchev, *Aster - agentic science toolkit for exoplanet research*, [2512.xxxxx](#).
- [46] “Genesis mission: A national mission to accelerate science through artificial intelligence.” <https://genesis.energy.gov/>.
- [47] “Monte carlo tools for beyond the standard model physics.” <https://theory.fnal.gov/mc4bsm/>.
- [48] I. Doršner, S. Fajfer, A. Greljo, J.F. Kamenik and N. Košnik, *Physics of leptoquarks in precision experiments and at particle colliders*, *Phys. Rept.* **641** (2016) 1 [[1603.04993](#)].
- [49] SHERPA collaboration, *Event generation with Sherpa 3*, *JHEP* **12** (2024) 156 [[2410.22148](#)].
- [50] A. Alloul, N.D. Christensen, C. Degrande, C. Duhr and B. Fuks, *FeynRules 2.0 - A complete toolbox for tree-level phenomenology*, *Comput. Phys. Commun.* **185** (2014) 2250 [[1310.1921](#)].

- [51] J. Alwall, R. Frederix, S. Frixione, V. Hirschi, F. Maltoni, O. Mattelaer et al., *The automated computation of tree-level and next-to-leading order differential cross sections, and their matching to parton shower simulations*, *JHEP* **07** (2014) 079 [[1405.0301](#)].
- [52] C. Bierlich et al., *A comprehensive guide to the physics and usage of PYTHIA 8.3*, *SciPost Phys. Codeb.* **2022** (2022) 8 [[2203.11601](#)].
- [53] M. Cacciari, G.P. Salam and G. Soyez, *FastJet User Manual*, *Eur. Phys. J. C* **72** (2012) 1896 [[1111.6097](#)].
- [54] I. Antcheva et al., *ROOT: A C++ framework for petabyte data storage, statistical analysis and visualization*, *Comput. Phys. Commun.* **180** (2009) 2499 [[1508.07749](#)].
- [55] A. Buckley, P. Ilten, D. Konstantinov, L. Lönnblad, J. Monk, W. Pokorski et al., *The HepMC3 event record library for Monte Carlo event generators*, *Comput. Phys. Commun.* **260** (2021) 107310 [[1912.08005](#)].
- [56] J. Alwall et al., *A Standard format for Les Houches event files*, *Comput. Phys. Commun.* **176** (2007) 300 [[hep-ph/0609017](#)].
- [57] J. Alwall, M. Herquet, F. Maltoni, O. Mattelaer and T. Stelzer, *MadGraph 5 : Going Beyond*, *JHEP* **06** (2011) 128 [[1106.0522](#)].
- [58] G. Bewick et al., *Herwig 7.3 release note*, *Eur. Phys. J. C* **84** (2024) 1053 [[2312.05175](#)].
- [59] DELPHES 3 collaboration, *DELPHES 3, A modular framework for fast simulation of a generic collider experiment*, *JHEP* **02** (2014) 057 [[1307.6346](#)].
- [60] J. Allison et al., *Recent developments in Geant4*, *Nucl. Instrum. Meth. A* **835** (2016) 186.
- [61] C. Bierlich et al., *Robust Independent Validation of Experiment and Theory: Rivet version 3*, *SciPost Phys.* **8** (2020) 026 [[1912.05451](#)].
- [62] OpenAI, :, S. Agarwal, L. Ahmad, J. Ai, S. Altman et al., *gpt-oss-120b and gpt-oss-20b model card*, [2508.10925](#).
- [63] Q. Jiang and G. Karniadakis, *Agenticsciml: Collaborative multi-agent systems for emergent discovery in scientific machine learning*, [2511.07262](#).
- [64] C. Andreopoulos et al., *The GENIE Neutrino Monte Carlo Generator*, *Nucl. Instrum. Meth. A* **614** (2010) 87 [[0905.2517](#)].
- [65] J. Isaacson, W.I. Jay, A. Lovato, P.A.N. Machado and N. Rocco, *Introducing a novel event generator for electron-nucleus and neutrino-nucleus scattering*, *Phys. Rev. D* **107** (2023) 033007 [[2205.06378](#)].
- [66] A. Alnuqaydan, S. Gleyzer and H. Prosper, *SYMBA: symbolic computation of squared amplitudes in high energy physics with machine learning*, *Mach. Learn. Sci. Tech.* **4** (2023) 015007 [[2206.08901](#)].

- [67] Z. Dong, K. Kong, K.T. Matchev and K. Matcheva, *Is the machine smarter than the theorist: Deriving formulas for particle kinematics with symbolic regression*, *Phys. Rev. D* **107** (2023) 055018 [[2211.08420](#)].
- [68] K.T. Matchev, K. Matcheva, P. Ramond and S. Verner, *Exploring the truth and beauty of theory landscapes with machine learning*, *Phys. Lett. B* **856** (2024) 138941 [[2401.11513](#)].
- [69] G.N. Wojcik, S.T. Eu and L.L. Everett, *Graph reinforcement learning for exploring model spaces beyond the standard model*, *Phys. Rev. D* **111** (2025) 035007 [[2407.07203](#)].
- [70] S. Kawai and N. Okada, *Truth, beauty, and goodness in grand unification: A machine learning approach*, *Phys. Lett. B* **860** (2025) 139221 [[2411.06718](#)].
- [71] Y.S. Koay, R. Enberg, S. Moretti and E. Camargo-Molina, *Generating particle physics Lagrangians with transformers*, [2501.09729](#).