

# TacitFlow: Learning Workflow Representations for Tacit-Knowledge-Grounded Machine Learning Engineering Agents

Yushan Jiang<sup>1</sup>, Wenchao Yu<sup>2†</sup>, Minyoung Choe<sup>3</sup>, Dongjin Song<sup>1†</sup>,  
Jingchao Ni<sup>4</sup>, Wei Cheng<sup>2</sup>, Haifeng Chen<sup>2</sup>

<sup>1</sup>University of Connecticut <sup>2</sup>NEC Laboratories America <sup>3</sup>KAIST <sup>4</sup>University of Houston

## Abstract

Large language model (LLM) agents have shown promise for automating machine learning engineering (MLE), but their iterative improvement remains weakly grounded. In particular, most existing methods lack structured representations of how expert pipeline components compose across tasks. Instead, their solution development is often driven by prior trials and flat textual context from external retrieval sources, with limited explicit support for reasoning over step dependencies and reusable pipeline compositions. To bridge this gap, we introduce TACITFLOW, a knowledge-grounded MLE agent built on structured representations of expert workflows. We distill expert solutions across image, tabular, and text domains into heterogeneous workflow graphs of canonicalized pipeline steps and execution dependencies. From this graph, we learn contrastive node embeddings that capture multi-modal step semantics from technical descriptions and code snippets, and train a task-conditioned ranking model to capture workflow performance preferences from competition leaderboard rankings. During runtime, these representations support a retrieve, edit, and rerank refinement loop, where structured workflow context guides explicit candidate edits and the task-conditioned ranking model prioritizes promising workflows before execution. This design avoids exhaustively running every candidate while making agent decisions traceable to individual workflow steps. Experiments across three modalities show that our workflow-structured design yields promising performance on a diverse subset of competitions from MLE-Bench, while case studies illustrate transparent, workflow-grounded agent behavior on end-to-end MLE tasks.

## CCS Concepts

• **Computing methodologies** → **Artificial intelligence; Machine learning.**

## Keywords

Representation Learning, Autonomous Data Science, Retrieval

## 1 Introduction

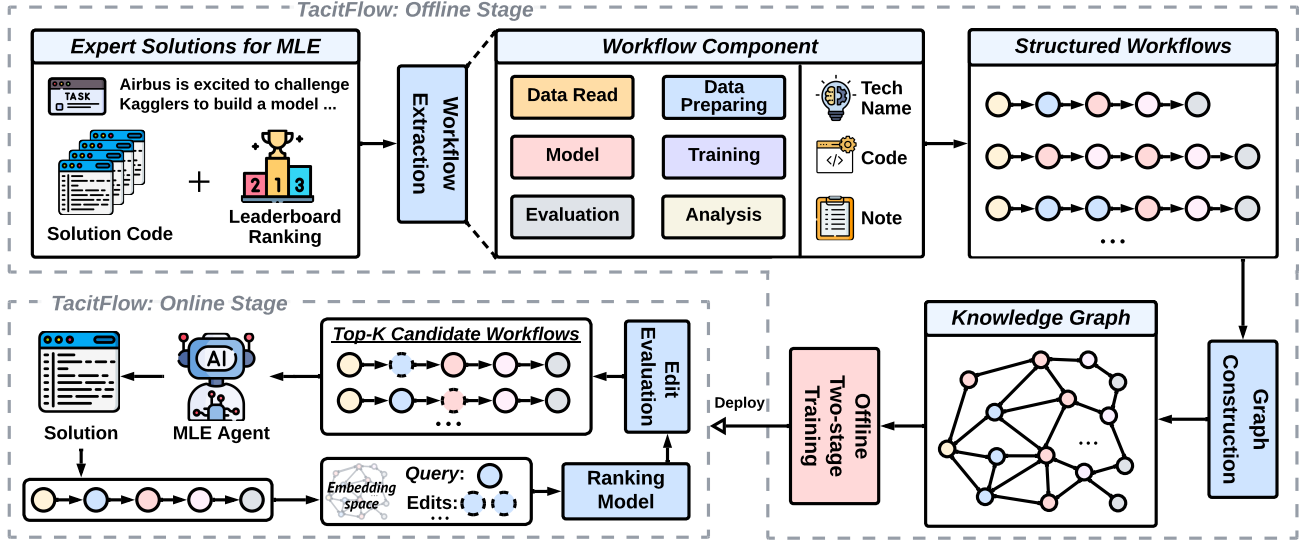
Building an effective machine learning solution for a new task is an iterative, knowledge-intensive process of exploring, evaluating, and adapting pipeline choices. Recent work on autonomous machine learning engineering (MLE) automates this process with LLM-based agents that propose, implement, and improve candidate solutions under a fixed budget [11, 26]. Despite promising progress, these agents continue to trail top human experts on real competitions [2], especially on tasks where strong solutions depend on tacit choices

that are non-obvious from the data alone, such as cleaning suspicious labels or result calibration. Such choices accumulate in public competition notebooks written by expert practitioners. A natural question to ask is, *how can we represent tacit expert knowledge as meaningful, actionable guidance for constructing and improving MLE solutions?*

Existing MLE agents have made solution development more systematic by organizing the exploration space of candidate solutions. Beyond basic trial-and-error, recent systems use tree-structured search, evolutionary exploration, cross-branch reuse, or hierarchical caches to make exploration more efficient and targeted [3, 4, 13, 14, 21, 23, 26]. These methods structure how agents explore and refine their own generated solutions, but they do not provide expert knowledge about which pipeline choices tend to work. Recent work has begun exploring retrieval-augmented agents that bring in external materials such as web pages, papers, or curated domain notes to expose the LLM to external expertise [5, 16, 18]. Yet the retrieved knowledge is mostly represented as raw text summaries, case examples, or high-level hints. As a result, expert knowledge may influence the LLM through context, but it rarely participates in the agent’s explicit decision process as a structured, reusable signal over pipelines that guides solution development.

To address this gap, we introduce TACITFLOW, an MLE framework that exploits tacit expert knowledge as structured workflow guidance. Rather than using expert knowledge as passive context for proposing solutions, TACITFLOW represents it as reusable pipeline steps, observed dependencies, and task-aware preferences over complete workflows. We first curate a heterogeneous workflow knowledge graph from 131 Kaggle competitions and 2,681 expert workflows across image, tabular, and text modalities. Each node corresponds to a canonical pipeline step with a technical name, natural-language notes, and code snippets, while each edge records an execution dependency observed in expert solutions, as shown in Figure 1. On top of this graph, we learn complementary signals at two granularities. At the node level, a self-supervised contrastive encoder produces step embeddings that capture both multi-modal semantic similarity and graph-structural proximity, yielding fine-grained step-level representations for retrieval. At the workflow level, a task-conditioned ranking model scores workflows by their expected quality on a given task. Supervised by competition leaderboard ranks, it provides a learned preference signal over candidate compositions. At runtime, TACITFLOW guides the agent’s next development step through a *retrieve-edit-rerank* loop. Given the target task and current solution state, the agent retrieves relevant steps via node embedding similarity, adapts them into candidate workflow changes, and uses the ranking model to select promising variants. The selected variants are then presented as suggested edits that condition the agent’s next proposal. This efficiently makes expert

<sup>†</sup> Wenchao Yu and Dongjin Song are corresponding authors.



**Figure 1: Overview of TACITFlow.** The offline stage extracts structured workflows from expert MLE solutions, canonicalizes them into a workflow knowledge graph, and learns node embeddings together with a task-conditioned ranking model. The online stage uses these learned guidance signals to map the agent’s current solution into a workflow, retrieve candidate edits, rerank edited workflows, and return selected workflow guidance to the MLE agent for refinement.

knowledge actionable without exhaustively testing every plausible modification. It also makes solution development more traceable, since each generated proposal can be linked to the retrieved steps attributed and prioritized by the rank. Our contributions are summarized as follows:

- We propose TACITFlow, a framework that turns tacit expert knowledge into a structured, learned representation for MLE agents, comprising reusable pipeline steps, execution dependencies, and task-aware quality preferences.
- We construct a heterogeneous workflow knowledge graph from 131 Kaggle competitions and 2,681 expert workflows across image, tabular, and text modalities, and propose bi-level representations over this graph to enable knowledge-grounded solution development.
- We integrate these learned representations into an iterative MLE agent through a retrieve-edit-rerank loop, making expert knowledge actionable and traceable while achieving performance gains on competitions across three modalities from MLE-Bench.

## 2 Related Work

**LLM agents for machine learning engineering.** Recent work on autonomous MLE formulates the problem as an iterative agent loop, with diverse systems advancing along two complementary axes: organizing the agent’s own exploration of candidate solutions, and grounding agents in external knowledge [10, 19, 21, 26]. On the exploration axis, AIDE [11] initiates the paradigm with greedy tree search, AIRA [21] enhances this with structured operator sets and scoped memory across exploration branches, and the dominant subsequent line adopts Monte Carlo Tree Search and its variants over candidate solutions [3, 4, 18], including extensions with cross-branch reference edges and budget-aware planning. MLE-STAR [16] explores along a single chain with ablation-driven refinement, while

another line pursues multi-agent evolutionary search across parallel populations [13, 23]. Memory mechanisms further enrich this axis: hierarchical cognitive caching of accumulated trials [14], and cross-branch trace reuse [3, 4] embed memory directly into the search loop. On the external knowledge axis, an early line of work uses curated case banks of expert solutions through case-based reasoning [7], while more recent agents bring expert materials into the LLM’s context through web search [16], curated knowledge bases of papers and tricks [18] as well as playbooks and code repositories [15], domain priors over models, data, and strategies [4], and library documentation paired with execution histories [5]. Despite this expanding access, existing systems consume external knowledge as raw text, case examples, or condensed prompts, without building expert solutions into a learned representation that captures pipeline steps, dependencies, and task-conditioned preferences. TACITFlow departs from this by introducing such a representation and integrating it into the agent’s proposal generation.

**Learned representations of workflows.** A growing line of work treats agentic workflows as graphs and learns representations over them. Workflow optimization methods [9, 27, 30] formulate workflow design as a search problem over communication topologies or code-represented workflows, typically optimized with reinforcement learning, meta-agent programming, or Monte Carlo Tree Search. Building on this paradigm, recent work has begun to learn binary performance predictors over agentic workflows so that promising candidates can be prioritized without expensive execution. FLORA-Bench [28] releases 600k workflow-task pairs over multi-agent communication topologies and trains graph-based predictors of workflow success on short reasoning and code-completion benchmarks. Furthermore, Agentic Predictor [22] adds code- and prompt-aware encoders, and GLOW [6] integrates LLM-based reasoning into graph predictors. We extend this line of work

into the substantially more demanding setting of MLE tasks. Compared to short reasoning or code-completion tasks solved by homogeneous agent communication, ML pipelines exhibit a heterogeneous and multi-stage compositional structure, including pre-processing, modeling, and training choices that interact in task-dependent ways, and competition leaderboards provide dense, fine-grained quality signals over complete workflows, going beyond binary success labels. Moreover, TACITFLOW shifts learned workflow representations from passive performance estimators to active components of an MLE agent’s decision process, reasoning over top candidate edits at each refinement step.

### 3 Problem Formulation

We consider the autonomous machine learning engineering task, where an agent produces a high-quality solution for a given ML competition under a fixed time budget. Formally, given an ML task with an associated dataset  $\mathcal{D} = \{\mathcal{D}_{\text{dev}}, \mathcal{D}_{\text{test}}\}$ , the agent iteratively develops candidate solutions  $s$  using only  $\mathcal{D}_{\text{dev}}$ , while final performance is measured by a task-specific metric  $M(s; \mathcal{D}_{\text{test}})$  on the held-out test set. Given a budget vector  $B$  (e.g., time and token usage), the agent’s goal is to solve the constrained optimization as:

$$s^* = \arg \max_{s \in \mathcal{S}} M(s; \mathcal{D}_{\text{test}}), \quad \text{s.t. } \text{budget}(s) \leq B \quad (1)$$

where  $\mathcal{S}$  is the feasible solution space and  $\text{time}(s)$  is the total elapsed time for producing  $s$ .

| Example                                                                                                     |                                                                                                      |
|-------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------|
| train_df = pd.read_csv('/kaggle/.../train.csv')<br>test_df = pd.read_csv('/kaggle/.../test.csv')            | <b>Data Read - data-read</b><br>Load train, test CSVs into dataframes                                |
| X = train_df.drop(columns=['target', 'id'])<br>X_test = test_df.drop(['id'], axis=1)                        | <b>Data Preparing - column-drop-id</b><br>Remove target columns from features                        |
| label_encoder = LabelEncoder()<br>y_encoded = label_encoder.fit_transform(y)                                | <b>Data Preparing - label-encoding</b><br>Encode labels into integer classes                         |
| scaler = StandardScaler()<br>X_scaled = scaler.fit_transform(X)<br>X_test_scaled = scaler.transform(X_test) | <b>Data Preparing - scaling-standard</b><br>Standardize features to zero mean / unit variance        |
| X_train, X_val, y_train, y_val = train_test_split(X_scaled, y_encoded, test_size=0.2)                       | <b>Data Preparing - stratified-split</b><br>Split features and labels into train and validation sets |
| model = XGBClassifier(n_estimators=100, ...)                                                                | <b>Modeling - xgbclassifier</b><br>Initialize XGBoost Classifier                                     |
| model.fit(X_train, y_train)                                                                                 | <b>Training - xgboost-train</b><br>Fit XGBClassifier on the training split                           |
| loss = log_loss(y_val, y_val_probs)                                                                         | <b>Evaluation - metric-log-loss</b><br>Compute multiclass log-loss on the validation set             |

Figure 2: An example of workflow extraction from a code solution.

### 4 Methodology

We propose a knowledge-grounded approach that equips the agent with structured workflow knowledge distilled from similar competition notebooks. For each data modality (e.g., image, tabular, text), we use an LLM-based extractor to convert top-performing solutions into evidential workflow steps that capture canonical operations, functional roles, and supporting code context, as shown in Figure 2. These steps are organized into a heterogeneous knowledge graph  $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ , where recurring operations across solutions are consolidated into shared nodes and observed execution dependencies define edges. The resulting graph exposes reusable workflow structure for step-level retrieval and workflow-level ranking. A complete

workflow  $w$  thus corresponds to a subgraph paired with a task description drawn from the competition metadata. The details are provided in Appendix B. Over this graph, we learn complementary signals at two granularities: node embeddings  $\mathbf{h}_v$  that capture multi-modal semantic similarity and graph-structural proximity, and a task-conditioned workflow ranking model that scores complete workflows by empirical quality. These bi-level representations serve as structured context for an iterative MLE agent, shifting solution development toward a knowledge-guided process that surfaces expert-derived candidates at every step and yields transparent decision traces.

#### 4.1 Learning Node Representations

**Multi-modal Graph Encoder.** We build a graph encoder  $f_\theta : \mathcal{G} \rightarrow \mathbb{R}^{|\mathcal{V}| \times d}$  with hidden dimension  $d$ . For each feature modality  $k \in \text{cat, tech, notes, code}$  (category label, technique name, technique notes, and code snippets), a projection  $\phi_k$  maps raw features to a modality-specific subspace with dimension  $d'_k$  chosen per modality (one-hot for cat, S-BERT [20] for tech, notes and CodeT5+ [24] for code). The projected features are concatenated and fused via a two-layer MLP:

$$\mathbf{h}_v^k = \phi_k(\mathbf{x}_v^k), \tilde{\mathbf{h}}_v = \text{MLP}([\mathbf{h}_v^{\text{cat}} \parallel \mathbf{h}_v^{\text{tech}} \parallel \mathbf{h}_v^{\text{notes}} \parallel \mathbf{h}_v^{\text{code}}])$$

The graph encoder consists of  $L$  GraphSAGE layers [8] and residual connections:

$$\mathbf{h}_v^{(\ell+1)} = \mathbf{h}_v^{(\ell)} + \text{GELU}\left(\text{SAGE}(\mathbf{h}_v^{(\ell)}, \{\mathbf{h}_u^{(\ell)}\}_{u \in \mathcal{N}(v)})\right)$$

Let  $\mathbf{h}_v = f_\theta(\mathcal{G})_v$  denote the encoder output after  $L$  layers. For contrastive pretraining, we attach a two-layer projection head  $g$  and obtain the contrastive embedding  $\mathbf{z}_v = g(\mathbf{h}_v)$ . Unless otherwise specified, we use  $\mathbf{h}_v$  (before  $g$ ) as the final node representation for downstream usage.

We pretrain node embeddings to distill these cross-workflow patterns together with each node’s multi-modal attributes via self-supervised contrastive learning, applying graph augmentations along edge perturbations and feature masking [29]. Specifically, two correlated graph views  $\tilde{\mathcal{G}}_1, \tilde{\mathcal{G}}_2$  are generated by independently applying (i) stochastic edge dropping with rate  $p_e$  and (ii) feature-dimension masking with rate  $p_f$ . The two views are processed by the shared encoder and projection head, yielding  $\mathbf{z}_v^1 = g(f_\theta(\tilde{\mathcal{G}}_1)_v)$  and  $\mathbf{z}_v^2 = g(f_\theta(\tilde{\mathcal{G}}_2)_v)$  with dimension  $d$ .

**Cross-view contrastive loss.** We first maximize agreement between representations of the same node across views via a symmetric InfoNCE objective [17]:

$$\mathcal{L}_{\text{cl}} = -\frac{1}{2N} \sum_{v \in \mathcal{V}} \sum_{a \neq b} \log \frac{\exp(\text{sim}(\mathbf{z}_v^a, \mathbf{z}_v^b)/\tau)}{\sum_{v' \in \mathcal{V}} \exp(\text{sim}(\mathbf{z}_v^a, \mathbf{z}_{v'}^b)/\tau)}$$

where  $a, b \in \{1, 2\}$  and  $a \neq b$  indicate two views,  $\text{sim}(\cdot, \cdot)$  is cosine similarity and  $\tau$  is a temperature parameter.

**VICReg regularization.** To prevent dimensional collapse of the learned representations [1], we add variance and covariance style regularization terms on both views. The variance term  $\mathcal{L}_{\text{var}}$  maintains per-dimension standard deviation, while the covariance term  $\mathcal{L}_{\text{cov}}$  penalizes off-diagonal entries of the covariance matrix  $C^i$  to

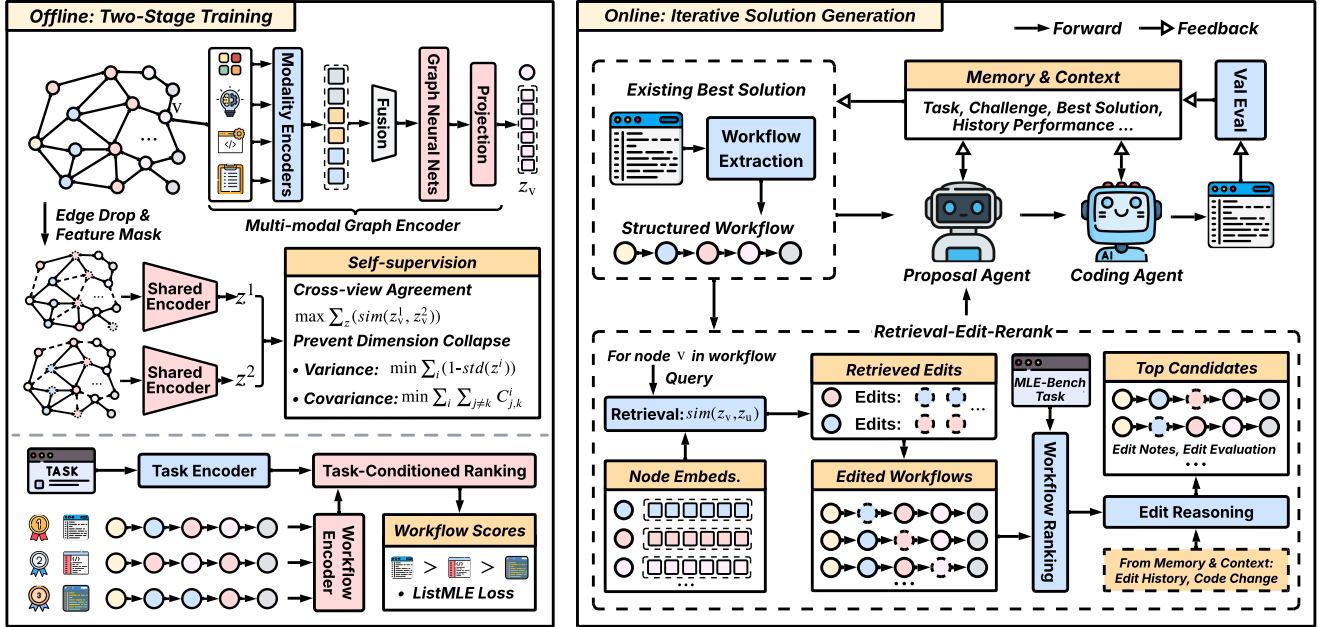


Figure 3: Methodology of TACITFLOW. Stage 1 learns multi-modal node embeddings over a heterogeneous workflow graph with self-supervised objectives. Stage 2 learns a task-conditioned ranking model over workflow representations to score complete workflows by expected quality. Stage 3 converts the current best solution into a structured workflow, retrieves candidate edits for a new task using node embeddings, reranks them with the ranking model, and refines the solution through a proposal agent under validation feedback and memory.

decorrelate dimensions:

$$\mathcal{L}_{\text{var}} = \frac{1}{2d} \sum_{i,j} \max(0, 1 - \text{std}(z_{:,j}^i)), \quad \mathcal{L}_{\text{cov}} = \frac{1}{2d} \sum_{i,j \neq k} (C_{jk}^i)^2$$

The full training objective is:

$$\mathcal{L} = \lambda_{\text{cl}} \mathcal{L}_{\text{cl}} + \lambda_{\text{var}} \mathcal{L}_{\text{var}} + \lambda_{\text{cov}} \mathcal{L}_{\text{cov}}. \quad (2)$$

## 4.2 Task-Conditioned Workflow Ranking

The node embeddings encode rich multi-modal semantics and graph-aware structural patterns, making them well suited for workflow level exploitation. To retain these properties, we fix the node embeddings and cast workflow selection as a ranking task within each competition.

Let  $\mathbf{X}_w \in \mathbb{R}^{n \times d}$  denote the sequence of node embeddings (padded to length  $n$ ) and  $\mathbf{M}_w \in \{0, 1\}^n$  the binary mask of valid positions. Task descriptions  $t_w$  (from the competition metadata) are encoded with Sentence-BERT and projected by an MLP as  $t'_w$ . In this stage, our goal is to learn a task-conditioned scoring function  $s_w$  that reflects the relative quality of workflows from the same competition. To this end, we employ a lightweight sequence encoder to obtain a workflow representation and then fuse it with the task embedding to predict a preference score.

**Workflow Encoder.** Because workflow steps execute sequentially, we encode the step sequence with a unidirectional LSTM that represents the execution dependencies:  $\mathbf{H}_w = \text{LSTM}(\mathbf{X}_w) \in \mathbb{R}^{n \times d_h}$ , where each hidden state  $\mathbf{h}_i^{\text{seq}}$  summarizes the workflow up to step  $i$ . To highlight which steps matter most, we apply an attentive pooling

mechanism that scores each hidden state against a single learned query vector  $\mathbf{q} \in \mathbb{R}^{d_h}$ :

$$\alpha_i = \frac{\mathbf{M}_{w,i} \exp(\mathbf{q}^\top \mathbf{h}_i^{\text{seq}})}{\sum_j \mathbf{M}_{w,j} \exp(\mathbf{q}^\top \mathbf{h}_j^{\text{seq}})}, \quad \mathbf{h}_w = \sum_i \alpha_i \mathbf{h}_i^{\text{seq}}$$

where the mask  $\mathbf{M}_w$  zeros out padded positions to handle varying workflow lengths. The attention weights  $\alpha_i$  serve as interpretable step-wise importance scores. We fuse the workflow and task representations via an MLP  $\phi_f$  and predict a scalar score as:  $s_w = \phi_f([\mathbf{t}'^\top \|\mathbf{h}_w])$ .

**Ranking Loss.** To supervise workflow ranking, we use the public leaderboard ranking of each competition as listwise supervision, yielding a ground-truth ranked list  $\sigma_c$  over the  $n_c$  workflows in competition  $c$ , ordered from best to worst. We adopt the ListMLE loss [25], a listwise objective that models the likelihood of the entire permutation under a Plackett-Luce model:

$$\mathcal{L}_{\text{rank}} = -\frac{1}{|C|} \sum_{c \in C} \sum_{j=1}^{n_c} \log \frac{\exp(s_{w_{\sigma_c(j)}})}{\sum_{\ell=j}^{n_c} \exp(s_{w_{\sigma_c(\ell)}})} \quad (3)$$

where  $C$  is the set of all competitions. This encourages the model to score workflows consistently with the ground-truth order across the entire candidate set.

## 4.3 TacitFlow-Informed MLE Agent

The workflow representations encode reusable knowledge about which solution patterns tend to perform well and how pipeline components relate structurally, which we seamlessly integrate into a solution development loop while maintaining transparent decision

traces. As our starting point, we adopt a base loop with components common to recent open-source MLE agents [11, 14, 18, 26]: tree-structured search that selectively expands promising branches, an exploration memory of past solutions and validation feedback, and stage-aware prompting strategies, and a standard coding workflow that implements proposals into runnable code via subset-based prototyping [26].

Within each search branch, we consider an MLE agent that alternates between proposing ideas, generating and evaluating code, and updating an exploration history  $\mathcal{H}_{<t}$ . Let  $\hat{c}_{t-1}$  be the latest accepted solution before iteration  $t$ , and  $\kappa_t$  be the structured context derived from workflow representations. The agent performs:

$$\begin{aligned} i_t &\leftarrow \text{PROPOSE}(\kappa_t, \mathcal{H}_{<t}) \\ c_t &\leftarrow \text{IMPLEMENT}(i_t; \hat{c}_{t-1}) \\ m_t &\leftarrow \text{EVALUATE}(c_t; \mathcal{D}_{\text{dev}}) \end{aligned}$$

where  $i_t$  is the initial/refinement proposal,  $c_t$  the resulting code, and  $m_t$  the validation performance. If  $c_t$  yields an improvement in  $m_t$  over the current baseline, we accept it as the new solution  $\hat{c}_t$ ; otherwise, we keep the baseline unchanged. The validation outcome and associated feedback, including execution status, error and debugging traces, and summaries of historical solutions, are appended to the history  $\mathcal{H}_{<t}$  to inform subsequent iterations.

Without external structure, an iterative MLE agent must propose refinements based solely on its parametric knowledge and accumulated trial history. While the history records past attempts and outcomes, it does not explicitly expose which workflow components are most promising to modify, nor does it provide a principled mechanism for reusing high-quality solution patterns observed in related tasks. To enable structured and traceable reasoning over solution components, we first convert the current solution code into the same workflow abstraction used in the offline stage:

$$w_{t-1} = \text{EXTRACT}(\hat{c}_{t-1}) = [v_0, v_1, \dots, v_n],$$

where EXTRACT denotes an LLM-based procedure that decomposes code into step-level components. This structured representation provides an interpretable view of the solution and allows subsequent refinement decisions to be attributed to specific steps.

**Step-level retrieval.** We first map each step in the current workflow  $w_{t-1}$  into the pretrained node embedding space. As the agent may produce steps that do not appear verbatim in the offline graph  $\mathcal{G}$ , we employ a three-tier matching scheme to bridge the gap between runtime workflows and the offline corpus: (i) Exact match: if the step’s standardized name appears in  $\mathcal{V}$ , we reuse the corresponding embedding  $\mathbf{z}_v$ . (ii) Fuzzy match: when no exact match exists, we check for substring containment, and use S-BERT cosine similarity, assigning the embedding of the closest matched node. (iii) Constructed embedding: when no sufficiently similar entry exists, we construct an embedding by feeding the step’s multi-modal features via the pretrained projection  $\phi_k$ . Given these embeddings, we retrieve semantically related alternatives for each step  $v \in \mathcal{V}_{w_{t-1}}$ :

$$\mathcal{N}_{\text{ref}}(v) = \{u \in \mathcal{V} : \text{sim}(\mathbf{z}_v, \mathbf{z}_u) \geq \delta, u \neq v\}$$

where  $\mathbf{z}_v$  denotes the pretrained node embedding. The retrieved nodes serve as candidate substitutions grounded in multi-modal semantic and structural similarity learned via graph pretraining.

**Edit and Rerank.** Each retrieved node  $u \in \mathcal{N}_{\text{ref}}(v)$  is treated as a candidate workflow edit. If  $u$  shares the same stage category as the original step  $v$ , it is used as an in-place substitute; otherwise, it suggests a cross-stage edit. Each candidate workflow  $\tilde{w}$  is generated by applying a single edit to  $w_{t-1}$ , and inherits the stage category of the edited step. We score each candidate with our task-conditioned workflow ranker  $s_t(\tilde{w})$ , and select the top candidates per stage category, preventing single-stage dominance and exposing the agent to refinements across the full pipeline:

$$\mathcal{W}_t^* = \bigcup_c \text{top-}k\{s_t(\tilde{w}) : \text{cat}(\tilde{w}) = c\}$$

The shortlist  $\mathcal{W}_t^*$  serves as the basis for the structured knowledge context  $\kappa_t$ , further enriched in the next step.

**Edit Reasoning and Proposal.**  $\mathcal{W}_t^*$  is task-conditioned but does not reflect the agent’s recent execution experience, namely, which edits have been tried, which succeeded or failed, and how the validation score has evolved. We therefore maintain a lightweight memory to inform reasoning over retrieved edits. At each iteration  $t$ , we enrich  $\mathcal{W}_t^*$  by conditioning each retrieved edit on the task context and accumulated lessons, attaching a short advisory note that identifies recent conflicts, suggests compatible edits, and highlights task-specific constraints. Informed by these notes and the current solution state, the proposal agent classifies each edit as adopt, consider, reject, or irrelevant, and organizes the resulting decisions and rationales into the structured knowledge context  $\kappa_t$  used to draft each proposal. It then drafts several candidate proposals, scores them, and either selects one or ensembles multiple candidates into the final proposal  $i_t$  passed to the coding agent. After execution, the proposal, code change, historical edits, and validation outcome are summarized and appended to memory, where they can inform edit reasoning in subsequent iterations. The complete procedure is summarized in Algorithm 1 (Appendix A).

## 5 Experiments

### 5.1 Experiment Setup

**Datasets, Baselines, and Metrics.** We evaluate TACITFLOW on a diverse 18-competition subset of MLE-Bench [2] spanning three input modalities (image, tabular, text), three difficulty tiers (easy, medium, hard), and diverse task types. For each competition we report the native evaluation score and the highest medal threshold attained (Above Median, Bronze, Silver, Gold), averaged over three runs. We compare against six open-source MLE agents spanning dominant design axes: AIDE [11], ML-Master [14], AIRA [21], R&D-Agent [26], Leeroo (KAPSO) [15], and MLEvolve. The baseline results are averaged from their submissions to the official MLE-Bench GitHub repository.

**Environment and LLM Configurations.** Each run uses a 12-hour wall-clock budget on a server with NVIDIA Quadro RTX 6000 (24 GB), 48 CPU cores, and 256 GB RAM. For runtime MLE we use gpt-5.1 with reasoning disabled (*none* effort). Offline workflow graph construction uses gpt-4.1 and gpt-4o-mini for workflow extraction, for canonicalization, and code feature cleansing.

**Workflow Graph.** The workflow graph supporting both offline representation learning and online edit retrieval is built from 131



**Table 1: Per-competition results on an 18-competition subset of MLE-bench. The mean of valid submission scores across that agent’s runs are reported. Medal coloring derived from the mean: ● gold, ● silver, ● bronze, ● above median, plain text otherwise. “–” = no valid submission across runs. \* indicates RD-Agent fallback results obtained with o3+GPT-4.1, used when the GPT-5 run produced no valid submission. ↑ higher is better, ↓ lower is better.**

| Task                      | Competition     | Metric             | AIDE       | ML-Master   | AIRA    | RD-Agent | Leeroo               | MLEvolve             | TacitFlow              |         |
|---------------------------|-----------------|--------------------|------------|-------------|---------|----------|----------------------|----------------------|------------------------|---------|
| Image                     | Classification  | dog-breed          | LogLoss↓   | 2.749       | ● 0.388 | 0.861    | ● 0.372 <sup>+</sup> | –                    | ● 0.379                | ● 0.364 |
|                           |                 | plant-pathology-20 | mROC-AUC↑  | ● 0.977     | ● 0.989 | ● 0.968  | ● 0.992              | ● 0.996              | ● 0.997                | ● 0.994 |
|                           |                 | dogs-vs-cats       | LogLoss↓   | 1.769       | ● 0.004 | 0.162    | ● 0.027              | ● 0.004              | ● 0.004                | ● 0.009 |
|                           | Segmentation    | tgs-salt           | mAP@IoU↑   | 0.394       | 0.252   | 0.643    | 0.793 <sup>+</sup>   | –                    | ● 0.836                | ● 0.824 |
|                           |                 | uw-gi-tract        | Dice+Haus↑ | 0.186       | 0.050   | 0.430    | 0.408 <sup>+</sup>   | –                    | 0.802                  | 0.571   |
|                           | Detection       | kuzushiji          | F1 Score↑  | 0.141       | 0.410   | 0.345    | 0.316                | ● 0.661              | ● 0.909                | ● 0.611 |
| Regression                | denoising-dirty | RMSE↓              | 0.082      | ● 0.011     | ● 0.047 | ● 0.011  | ● 0.031              | ● 0.008              | ● 0.011                |         |
| Text                      | Classification  | spooky-author      | LogLoss↓   | 0.489       | ● 0.279 | ● 0.247  | ● 0.259              | ● 0.254              | ● 0.214                | ● 0.301 |
|                           |                 | detect-insults     | AUC↑       | ● 0.905     | ● 0.922 | ● 0.949  | ● 0.928              | ● 0.904              | ● 0.954                | ● 0.945 |
|                           |                 | jigsaw-toxic       | mROC-AUC↑  | 0.970       | ● 0.983 | ● 0.986  | ● 0.986              | ● 0.987              | ● 0.988                | ● 0.984 |
|                           | Regression      | us-patent          | Pearson↑   | 0.718       | 0.849   | ● 0.860  | 0.801                | ● 0.870              | ● 0.865                | 0.818   |
|                           |                 | essay-score-2      | QWK↑       | 0.737       | 0.821   | 0.826    | 0.825                | ● 0.859              | ● 0.838                | ● 0.830 |
|                           | Seq2Seq         | text-norm-russian  | Accuracy↑  | 0.927       | 0.963   | 0.968    | ● 0.981              | ● 0.976              | ● 0.989                | ● 0.978 |
| Tabular                   | Classification  | random-acts        | AUC↑       | ● 0.643     | ● 0.607 | ● 0.629  | ● 0.786              | –                    | ● 0.700                | ● 0.797 |
|                           |                 | tps-dec-2021       | Accuracy↑  | 0.928       | ● 0.962 | ● 0.954  | ● 0.963              | ● 0.963              | ● 0.962                | ● 0.963 |
|                           | Regression      | nomad2018          | mRMSLE↓    | 0.102       | ● 0.063 | 0.085    | ● 0.060              | ● 0.062              | ● 0.059                | ● 0.059 |
|                           |                 | tps-may-2022       | AUC↑       | 0.914       | ● 0.992 | ● 0.990  | 0.970 <sup>+</sup>   | –                    | ● 0.998                | ● 0.979 |
|                           |                 | stanford-vaccine   | MCRMSE↓    | 0.465       | ● 0.247 | 0.461    | ● 0.269              | ● 0.333              | ● 0.214                | ● 0.236 |
| LLM Model & Configuration |                 |                    | o1-preview | deepseek-r1 | o3      | GPT-5    | Gemini-3-Pro-Preview | Gemini-3-Pro-Preview | GPT-5.1 (No Reasoning) |         |
| Execution Time (Hours)    |                 |                    | 24         | 12          | 24      | 12       | 24                   | 12                   | 12                     |         |

**Table 2: Summary statistics of the collected workflow corpus.**

| Statistic                    | Image  | Tabular | Text  |
|------------------------------|--------|---------|-------|
| Competitions                 | 48     | 55      | 28    |
| Workflows                    | 883    | 1,226   | 572   |
| Nodes                        | 3,536  | 3,836   | 2,754 |
| Edges                        | 10,356 | 14,510  | 8,054 |
| Avg. edges / node            | 2.93   | 3.78    | 2.92  |
| Mean workflows / competition | 18.4   | 22.3    | 20.4  |

historical Kaggle competitions and 2,681 expert workflows. Summary statistics of the corpus, including modality and stage-category coverage, average workflow length, and node degree distributions are provided in Table 2. Appendix B provides the list of competitions and details of workflow collection, extraction, graph construction, canonicalization, deduplication, and code quality control.

**Leakage Control.** Before crawling notebooks, we removed all competitions appearing in MLE-Bench using official competition IDs and manually checked aliases. During evaluation, the agent has no internet access and cannot retrieve target-competition notebooks, discussions, or solutions. The workflow graph and ranker supervision are constructed only from competitions disjoint from the 18 evaluation competitions listed in Appendix Table 8.

**Hyperparameters.** For offline training, we use Optuna for hyperparameter selection, using Adam [12] as the optimizer. For online MLE, we use two branches in Monte Carlo Tree Search, and select up to 3 top candidates per category in our retrieve-edit-rerank loop.

## 5.2 Runtime Experiments

Table 1 reports per-competition results on the 18-competition subset of MLE-Bench. TACITFlow uses gpt-5.1 with reasoning effort set to *none*, whereas most baselines rely on reasoning models such as o1-preview, DeepSeek-R1, o3, or Gemini-3-Pro-Preview. Despite using this non-reasoning configuration, TACITFlow achieves the best score on 4 competitions and the second-best score on 5, reaches at least above-median performance on 16 of 18 competitions, and produces valid submissions for all 18. Among matched-budget systems with a 12-hour wall-clock limit, TACITFlow matches or exceeds R&D-Agent (GPT-5) on a majority of competitions. In contrast, Leeroo runs with a 24-hour budget using Gemini-3-Pro-Preview but fails to produce valid submissions on 5 of 18 competitions, including 4 medium-complexity tasks (tgs-salt, uw-gi-tract, random-acts, and tps-may-2022). This highlights a robustness gap beyond aggregated medal rates. MLEvolve is the best baseline in the current MLE-bench open-source leaderboard. The gap to TACITFlow concentrates on tasks demanding specialized architectural exploration, notably segmentation and detection (uw-gi-tract, kuzushiji) where MLEvolve’s reasoning model is most beneficial. On tabular tasks, where solution quality is dominated by reusable feature engineering and modeling patterns rather than open-ended reasoning, TACITFlow reaches the best score on 3 of 4 competitions (random-acts, tps-dec-2021, nomad2018), matching or surpassing MLEvolve despite using a non-reasoning LLM. We read this as evidence that explicit, learned workflow representations are a

complementary signal that can recover much of the gap from a non-reasoning backbone.

### 5.3 Ablation Study

**The effectiveness of the retrieve-edit-rerank design.** Table 3 evaluates the retrieve-edit-rerank design on five medium- and hard-complexity MLE-Bench competitions. We compare TacitFlow with a *Vanilla* variant that uses the same iterative MLE scaffold but disables the workflow knowledge graph, and report both early progress and final performance. Across these challenging tasks, TacitFlow improves final performance on all five competitions, with an average improvement of 45.4%. This suggests that retrieved expert workflow components provide useful guidance when strong solutions require nontrivial pipeline choices. The same design also improves early-stage robustness: TacitFlow reaches the first valid submission faster on four of the five competitions and reduces the average first-valid time from 6.86 to 5.73 hours, corresponding to a 14.1% improvement. These gains indicate that retrieve-edit-rerank benefits both dimensions of the agent loop: it helps the agent reach executable solutions sooner and steers subsequent refinement toward higher-performing workflow variants, rather than spending the budget rediscovering common pipeline patterns through less-focused trial and error.

**Table 3: Effect of the retrieve-edit-rerank design on early progress and final performance across five medium and hard competitions. First-valid denotes the time to obtain the first valid submission.**

| Competition         | First-valid Time (h) |         |          | Performance        |                    |          |
|---------------------|----------------------|---------|----------|--------------------|--------------------|----------|
|                     | Ours                 | Vanilla | $\Delta$ | Ours               | Vanilla            | $\Delta$ |
| kuzushiji           | 7.92                 | 9.34    | -1.42    | 0.611              | 0.220              | +0.391   |
| tgs-salt            | 4.29                 | 4.66    | -0.37    | 0.824              | 0.761              | +0.063   |
| uw-gi-tract         | 5.65                 | 6.41    | -0.76    | 0.523              | 0.416              | +0.107   |
| stanford-vaccine    | 3.91                 | 4.18    | -0.17    | $\downarrow$ 0.236 | $\downarrow$ 0.268 | +0.032   |
| essay-score-2       | 6.88                 | 9.72    | -2.84    | 0.830              | 0.802              | +0.028   |
| <b>Average Imp.</b> | 14.1% (5.73 vs 6.86) |         |          | 45.4%              |                    |          |

**Component ablation of learned guidance.** We further ablate the representation learning and ranking components that provide learned guidance in TacitFlow. We use two complementary evaluations. The first is an in-corpus evaluation, which performs cross-validation over the expert workflow corpus used to build the workflow graph. The second is an out-of-corpus evaluation on solution trajectories from our TacitFlow MLE-Bench runs, where each intermediate solution is converted into the same workflow abstraction and ranked by the learned model. In this setting, the ground-truth rank is determined by the official MLE-Bench score of each submission, providing a direct test of whether the learned model captures true competition-level workflow quality. Tables 4 and 5 evaluate whether the learned guidance model recovers the true quality ordering of workflows. We report NDCG@5, MRR, and Top-3 accuracy to measure top-ranked workflow quality, the rank of the best workflow, and whether it appears among the top candidates. Across both in-corpus and out-of-corpus evaluations, the full model achieves the strongest overall performance among all variants, showing

**Table 4: Model ablation (in-corpus, 5-fold CV  $\times$  3 seeds).**

|         | Condition     | NDCG@5                              | MRR                                 | Top-3                               |
|---------|---------------|-------------------------------------|-------------------------------------|-------------------------------------|
| Image   | w/o Pretrain. | 0.888 $\pm$ 0.032                   | 0.449 $\pm$ 0.133                   | 0.584 $\pm$ 0.195                   |
|         | w/o VICReg.   | 0.880 $\pm$ 0.026                   | 0.450 $\pm$ 0.118                   | 0.574 $\pm$ 0.128                   |
|         | w/o Attn.     | <u>0.889 <math>\pm</math> 0.022</u> | <u>0.459 <math>\pm</math> 0.122</u> | 0.575 $\pm$ 0.159                   |
|         | Ours          | <b>0.897 <math>\pm</math> 0.037</b> | <b>0.488 <math>\pm</math> 0.132</b> | <b>0.612 <math>\pm</math> 0.180</b> |
| Text    | w/o Pretrain. | 0.786 $\pm$ 0.061                   | 0.315 $\pm$ 0.110                   | 0.324 $\pm$ 0.181                   |
|         | w/o VICReg.   | 0.804 $\pm$ 0.067                   | <u>0.427 <math>\pm</math> 0.073</u> | 0.522 $\pm$ 0.135                   |
|         | w/o Attn.     | <u>0.826 <math>\pm</math> 0.060</u> | 0.407 $\pm$ 0.104                   | <b>0.547 <math>\pm</math> 0.181</b> |
|         | Ours          | <b>0.833 <math>\pm</math> 0.051</b> | <b>0.434 <math>\pm</math> 0.139</b> | <u>0.529 <math>\pm</math> 0.258</u> |
| Tabular | w/o Pretrain. | 0.775 $\pm$ 0.093                   | 0.493 $\pm$ 0.150                   | 0.673 $\pm$ 0.231                   |
|         | w/o VICReg.   | <u>0.795 <math>\pm</math> 0.055</u> | <u>0.534 <math>\pm</math> 0.113</u> | <u>0.709 <math>\pm</math> 0.204</u> |
|         | w/o Attn.     | 0.769 $\pm$ 0.079                   | 0.475 $\pm$ 0.166                   | 0.636 $\pm$ 0.220                   |
|         | Ours          | <b>0.814 <math>\pm</math> 0.058</b> | <b>0.540 <math>\pm</math> 0.109</b> | <b>0.720 <math>\pm</math> 0.215</b> |

**Table 5: Model ablation (out-of-corpus evaluation on MLE-Bench runs, 3 seeds).**

|         | Condition     | NDCG@5                              | MRR                                 | Top-3                               |
|---------|---------------|-------------------------------------|-------------------------------------|-------------------------------------|
| Image   | w/o Pretrain. | 0.645 $\pm$ 0.019                   | <u>0.258 <math>\pm</math> 0.012</u> | 0.202 $\pm$ 0.110                   |
|         | w/o VICReg.   | 0.650 $\pm$ 0.002                   | 0.256 $\pm$ 0.008                   | 0.260 $\pm$ 0.015                   |
|         | w/o Attn.     | <u>0.675 <math>\pm</math> 0.022</u> | 0.242 $\pm$ 0.008                   | <u>0.271 <math>\pm</math> 0.059</u> |
|         | Ours          | <b>0.688 <math>\pm</math> 0.016</b> | <b>0.284 <math>\pm</math> 0.025</b> | <b>0.344 <math>\pm</math> 0.026</b> |
| Text    | w/o Pretrain. | <u>0.767 <math>\pm</math> 0.021</u> | 0.315 $\pm$ 0.025                   | <u>0.508 <math>\pm</math> 0.011</u> |
|         | w/o VICReg.   | 0.746 $\pm$ 0.025                   | <u>0.391 <math>\pm</math> 0.020</u> | 0.413 $\pm$ 0.040                   |
|         | w/o Attn.     | 0.724 $\pm$ 0.017                   | 0.389 $\pm$ 0.028                   | 0.421 $\pm$ 0.045                   |
|         | Ours          | <b>0.783 <math>\pm</math> 0.023</b> | <b>0.519 <math>\pm</math> 0.008</b> | <b>0.587 <math>\pm</math> 0.092</b> |
| Tabular | w/o Pretrain. | 0.561 $\pm$ 0.013                   | 0.178 $\pm$ 0.020                   | 0.083 $\pm$ 0.059                   |
|         | w/o VICReg.   | 0.567 $\pm$ 0.042                   | 0.194 $\pm$ 0.012                   | 0.208 $\pm$ 0.059                   |
|         | w/o Attn.     | <u>0.638 <math>\pm</math> 0.019</u> | <u>0.281 <math>\pm</math> 0.027</u> | <b>0.375 <math>\pm</math> 0.000</b> |
|         | Ours          | <b>0.650 <math>\pm</math> 0.024</b> | <b>0.433 <math>\pm</math> 0.023</b> | <b>0.375 <math>\pm</math> 0.000</b> |

that the learned representations provide a useful workflow preference signal. Contrastive pretraining with variance and covariance regularization brings consistent gains, suggesting the importance of robust self-supervised step representations that capture both semantic similarity and structural neighborhoods. Step-level attention further improves ranking by adaptively weighting workflow components according to their relevance to solution quality.

On the other hand, the transition from in-corpus evaluation to MLE-Bench trajectory evaluation introduces a moderate performance drop, reflecting the distribution shift from curated expert workflows to unseen agent-generated solutions. Nevertheless, the learned workflow guidance remains effective under this harder out-of-corpus setting, capturing useful ranking signals that generalize beyond the expert workflow corpus. We provide the full results with more metrics (NDCG@10, Top1) in Table 12 and 13, as well as other ablations including the feature importance (technique name, code, notes) across three image, text, and tabular modalities, in Table 9, 10, 11 of Appendix B.5.

### 5.4 Case Study: Traceable Refinement

We provide a case study that illustrates a single refinement step on plant-pathology-2020-fgvc7, as shown in Figure 4. Starting from the agent’s current solution, TacitFlow first extracts a 21-step workflow spanning data preparation, modeling, training, and

| Extracted Workflow and Attention Score |                  |                          |             | Final Proposal and Edit Guidance                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                    |
|----------------------------------------|------------------|--------------------------|-------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Step 1                                 | [Data Read]      | data-read                | exact       | 0.003                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 2                                 | [Data Preparing] | label-encoding           | exact       | 0.052                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 3                                 | [Data Preparing] | stratified-k-fold        | partial     | 0.013                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 4                                 | [Data Preparing] | resize                   | partial     | 0.007                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 5                                 | [Data Preparing] | random-resized-crop      | exact       | 0.088                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 6                                 | [Data Preparing] | random-horizontal-flip   | partial     | 0.232                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 7                                 | [Data Preparing] | color-jitter             | exact       | 0.106                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 8                                 | [Data Preparing] | gaussian-blur            | exact       | 0.060                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 9                                 | [Data Preparing] | to-tensor                | partial     | 0.008                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 10                                | [Data Preparing] | normalization            | partial     | 0.090                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 11                                | [Data Preparing] | random-erasing           | exact       | 0.036                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 12                                | [Data Preparing] | mixup                    | partial     | 0.006                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 13                                | [Data Preparing] | class-weight-computation | constructed | 0.030                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 14                                | [Modeling]       | efficientnet-b2          | exact       | 0.192                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 15                                | [Training]       | adamw                    | partial     | 0.019                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 16                                | [Training]       | cosine-annealing         | partial     | 0.009                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 17                                | [Training]       | grad-clip                | constructed | 0.028                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 18                                | [Training]       | label-smoothing          | exact       | 0.014                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 19                                | [Training]       | focal-loss               | exact       | 0.004                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 20                                | [Training]       | early-stopping           | exact       | 0.001                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Step 21                                | [Evaluation]     | roc-auc                  | exact       | 0.003                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|                                        |                  |                          |             | <b>Final Proposal:</b> Swap the EfficientNet-B2 backbone for an ImageNet-pretrained EfficientNet-B3 at a slightly higher resolution while keeping the rest of the pipeline identical: use torchvision's efficientnet_b3 with IMAGENET1K_V1 weights, set input_size=320 and crop_size=300 in the config (and corresponding Resize/CenterCrop in transforms), retain the same 4-way softmax head, soft-label focal focal loss with class weights, MixUp, two-stage fine-tuning with gradual unfreezing, 5-fold stratified CV, and early stopping. This modestly larger, higher-resolution backbone should better capture fine-grained disease patterns on leaves and yield a small but meaningful improvement in mean column-wise ROC AUC within the existing ~1-hour runtime budget. |
|                                        |                  |                          |             | <b>Edit Candidates - Modeling:</b><br><b>Option 1:</b> Replace efficientnet-b2 → efficientnetb3-imagenet-pretrained<br><b>Similarity Score: 0.684      Ranking Score <math>\Delta</math> +0.2289</b><br><b>Reasoning: Adopted</b> - This directly follows the suggested EfficientNetB3 backbone swap with minimal recipe changes, aligned with prior success from B0→B2.                                                                                                                                                                                                                                                                                                                                                                                                            |
|                                        |                  |                          |             | <b>Option 2:</b> Replace efficientnet-b2 → resnet50v2-imagenet-pretrained<br><b>Similarity Score: 0.621      Ranking Score <math>\Delta</math> -0.1329</b><br><b>Reasoning: Considered</b> - ResNet50V2 is a plausible second backbone for diversity, but starting with B2+B3 keeps preprocessing consistent and simplifies integration.                                                                                                                                                                                                                                                                                                                                                                                                                                            |

**Figure 4: Case study on *plant-pathology-2020-fgvc7*. Left: the current solution is represented as an extracted workflow with step-level attention scores. Right: TacitFlow retrieves and ranks candidate modeling edits, then forms the final proposal by adopting the EfficientNet-B3 backbone upgrade.**

evaluation. Each step is represented by a stage category, canonical technique name, matching status, and an attention score from the workflow ranker. The matching status records how the runtime step is connected to the offline workflow graph: common operations such as data-read, label-encoding, and efficientnet-b2 are matched exactly, while several augmentation and training steps are partially matched or constructed. This gives the agent an explicit intermediate representation of the current solution, rather than treating the code as a less clear prompt context. The step-level attention scores further expose which components the ranking model considers most consequential under the current task and workflow context. The largest attention weights identify the workflow components that the ranker considers most performance-sensitive under the current task context, including data augmentation, normalization, and the efficientnet-b2 backbone. This pattern separates routine implementation components from the parts of the workflow that are likely to affect final performance. In this example, the ranking model identifies the modeling backbone and the surrounding augmentation recipe as the relevant targets for refinement.

In the right panel, we provide a representative edit for the attended modeling step. Given the current efficientnet-b2 component, TACITFLOW retrieves semantically related backbone substitutions from the workflow graph and reranks them under the current workflow context. The reranked edit candidates include replacing efficientnet-b2 with a pretrained efficientnet-b3 or resnet50v2, which reflects different preferences. Conditioned on the task context and accumulated memory, the proposal agent adopts the efficientnet-b3 replacement as a minimal, recipe-consistent upgrade, since it builds on the current successful EfficientNet pipeline while increasing capacity and input resolution. Overall, the final edit is supported by a traceable evidence chain:

workflow extraction identifies the relevant pipeline steps, attention highlights the performance-sensitive modeling and augmentation components, retrieval surfaces alternative choices, reranking estimates their compatibility with the current workflow, and the proposal agent instantiates the adopted edit with a rationale. This allows the refinement to be attributed to concrete workflow components and alternatives, rather than appearing as an unsupported free-form suggestion. The full demonstration for this competition is provided in Appendix D.

## 6 Conclusion

We introduced TACITFLOW, a framework for turning tacit expert ML knowledge into structured, learned workflow guidance for iterative MLE agents. By constructing a heterogeneous workflow knowledge graph from 131 competitions and 2,681 expert workflows, TACITFLOW learns step-level representations and task-conditioned workflow preferences that support a retrieve-edit-rerank loop during solution development. On a diverse 18-competition subset of MLE-Bench, TACITFLOW improves automated MLE performance while making each refinement traceable to concrete workflow steps, retrieved edits, and ranking signals. Future work includes extending the framework to additional modalities such as time series and signal data, integrating the learned representations with stronger search frameworks and frontier reasoning models, and scaling the workflow corpus beyond Kaggle.

## Acknowledgments

Dongjin Song and Yushan Jiang gratefully acknowledge the support of the National Science Foundation under Grant No. 2338878, as well as the generous research gift from NEC Labs America.



## References

- [1] Adrien Bardes, Jean Ponce, and Yann LeCun. [n. d.]. VICReg: Variance-Invariance-Covariance Regularization for Self-Supervised Learning. In *International Conference on Learning Representations*.
- [2] Jun Shern Chan, Neil Chowdhury, Oliver Jaffe, James Aung, Dane Sherburn, Evan Mays, Giulio Starace, Kevin Liu, Leon Maksin, Tejal Patwardhan, et al. [n. d.]. MLE-bench: Evaluating Machine Learning Agents on Machine Learning Engineering. In *The Thirteenth International Conference on Learning Representations*.
- [3] Jiefeng Chen, Bhavana Dalvi Mishra, Jaehyun Nam, Rui Meng, Tomas Pfister, and Jinsung Yoon. 2026. MARS: Modular Agent with Reflective Search for Automated AI Research. *arXiv preprint arXiv:2602.02660* (2026).
- [4] Shangheng Du, Xiangchao Yan, Dengyang Jiang, Jiakang Yuan, Yusong Hu, Xin Li, Liang He, Bo Zhang, and Lei Bai. 2025. AutoMLGen: Navigating Fine-Grained Optimization for Coding Agents. *arXiv preprint arXiv:2510.08511* (2025).
- [5] Haoyang Fang, Boran Han, Nick Erickson, Xiyuan Zhang, Su Zhou, Anirudh Daggar, Jiani Zhang, Ali Caner Turkmen, Cuixiong Hu, Huzefa Rangwala, et al. [n. d.]. MLZero: A Multi-Agent System for End-to-end Machine Learning Automation. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*.
- [6] Wei Guan, Jian Cao, Jinyu Cai, Qiqi Cai, Jianqi Gao, and See-Kiong Ng. 2025. GLOW: Graph-Language Co-Reasoning for Agentic Workflow Performance Prediction. *arXiv preprint arXiv:2512.15751* (2025).
- [7] Siyuan Guo, Cheng Deng, Ying Wen, Hechang Chen, Yi Chang, and Jun Wang. 2024. DS-Agent: Automated Data Science by Empowering Large Language Models with Case-Based Reasoning. In *International Conference on Machine Learning*. PMLR, 16813–16848.
- [8] Will Hamilton, Zhitao Ying, and Jure Leskovec. 2017. Inductive representation learning on large graphs. *Advances in neural information processing systems* 30 (2017).
- [9] Shengran Hu, Cong Lu, and Jeff Clune. [n. d.]. Automated Design of Agentic Systems. In *The Thirteenth International Conference on Learning Representations*.
- [10] Qian Huang, Jian Vora, Percy Liang, and Jure Leskovec. 2024. MLAGentBench: Evaluating Language Agents on Machine Learning Experimentation. In *International Conference on Machine Learning*. PMLR, 20271–20309.
- [11] Zhengyao Jiang, Dominik Schmidt, Dhruv Srikanth, Dixing Xu, Ian Kaplan, Deniss Jacenko, and Yuxiang Wu. 2025. Aide: Ai-driven exploration in the space of code. *arXiv preprint arXiv:2502.13138* (2025).
- [12] Diederik P Kingma and Jimmy Ba. 2014. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980* (2014).
- [13] Annan Li, Chufan Wu, Zengle Ge, Yee Hin Chong, Zhihan Hou, Lizhe Cao, Cheng Ju, Jianmin Wu, Huaiming Li, Haobo Zhang, et al. 2025. The fm agent. *arXiv preprint arXiv:2510.26144* (2025).
- [14] Zexi Liu, Yuzhu Cai, Xinyu Zhu, Yujie Zheng, Runkun Chen, Ying Wen, Yanfeng Wang, Siheng Chen, et al. 2025. MI-master: Towards ai-for-ai via integration of exploration and reasoning. *arXiv preprint arXiv:2506.16499* (2025).
- [15] Alireza Nadafian, Alireza Mohammadshahi, and Majid Yazdani. 2026. KAPSO: A Knowledge-grounded framework for Autonomous Program Synthesis and Optimization. *arXiv preprint arXiv:2601.21526* (2026).
- [16] Jaehyun Nam, Jinsung Yoon, Jiefeng Chen, Jinwoo Shin, Sercan O Arik, and Tomas Pfister. [n. d.]. MLE-STAR: Machine Learning Engineering Agent via Search and Targeted Refinement. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*.
- [17] Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748* (2018).
- [18] Yixin Ou, Yujie Luo, Jingsheng Zheng, Lanning Wei, Zhuoyun Yu, Shuofei Qiao, Jintian Zhang, Da Zheng, Yuren Mao, Yunjun Gao, et al. 2025. Automind: Adaptive knowledgeable agent for automated data science. *arXiv preprint arXiv:2506.10974* (2025).
- [19] Rushi Qiang, Yuchen Zhuang, Yinghao Li, Dingu Sagar VK, Rongzhi Zhang, Changhao Li, Ian Shu-Hei Wong, Sherry Yang, Percy Liang, Chao Zhang, et al. [n. d.]. MLE-Dojo: Interactive Environments for Empowering LLM Agents in Machine Learning Engineering. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems Datasets and Benchmarks Track*.
- [20] Nils Reimers and Iryna Gurevych. 2019. Sentence-bert: Sentence embeddings using siamese bert-networks. In *Proceedings of the 2019 conference on empirical methods in natural language processing and the 9th international joint conference on natural language processing (EMNLP-IJCNLP)*. 3982–3992.
- [21] Edan Toledo, Karen Hambardzumyan, Martin Josifoski, Rishi Hazra, Nicolas Baldwin, Alexis Audran-Reiss, Michael Kuchnik, Despoina Magka, Minqi Jiang, Alisia Maria Lupidi, et al. 2025. Ai research agents for machine learning: Search, exploration, and generalization in mle-bench. *arXiv preprint arXiv:2507.02554* (2025).
- [22] Patara Tirat, Wonyong Jeong, and Sung Ju Hwang. 2026. Multi-View Encoders for Performance Prediction in LLM-Based Agentic Workflows. In *The Fourteenth International Conference on Learning Representations*. <https://openreview.net/forum?id=7oeKDZsmWp>
- [23] Chunhui Wan, Xunan Dai, Zhuo Wang, Minglei Li, Yanpeng Wang, Yinan Mao, Yu Lan, and Zhiwen Xiao. 2025. Loongflow: Directed evolutionary search via a cognitive plan-execute-summarize paradigm. *arXiv preprint arXiv:2512.24077* (2025).
- [24] Yue Wang, Hung Le, Akhilesh Gotmare, Nghi Bui, Junnan Li, and Steven Hoi. 2023. Codet5+: Open code large language models for code understanding and generation. In *Proceedings of the 2023 conference on empirical methods in natural language processing*. 1069–1088.
- [25] Fen Xia, Tie-Yan Liu, Jue Wang, Wensheng Zhang, and Hang Li. 2008. Listwise approach to learning to rank: theory and algorithm. In *Proceedings of the 25th international conference on Machine learning*. 1192–1199.
- [26] Xu Yang, Xiao Yang, Shikai Fang, Yifei Zhang, Jian Wang, Bowen Xian, Qizheng Li, Jingyuan Li, Minrui Xu, Yuante Li, et al. 2025. R&D-Agent: An LLM-Agent Framework Towards Autonomous Data Science. *arXiv preprint arXiv:2505.14738* (2025).
- [27] Jiayi Zhang, Jinyu Xiang, Zhaoyang Yu, Fengwei Teng, Xiong-Hui Chen, Jiaqi Chen, Mingchen Zhuge, Xin Cheng, Sirui Hong, Jinlin Wang, et al. [n. d.]. AFLOW: Automating Agentic Workflow Generation. In *The Thirteenth International Conference on Learning Representations*.
- [28] Yuanshuo Zhang, Yuchen Hou, Bohan Tang, Shuo Chen, Muhan Zhang, Xiaowen Dong, and Siheng Chen. 2025. FLORA: GNNs as Predictors of Agentic Workflow Performances. In *The Fourth Learning on Graphs Conference*. <https://openreview.net/forum?id=RmLDI5kUxa>
- [29] Yanqiao Zhu, Yichen Xu, Feng Yu, Qiang Liu, Shu Wu, and Liang Wang. 2020. Deep graph contrastive representation learning. *arXiv preprint arXiv:2006.04131* (2020).
- [30] Mingchen Zhuge, Wenyi Wang, Louis Kirsch, Francesco Faccio, Dmitrii Khizbullin, and Jürgen Schmidhuber. 2024. GPTSwarm: Language Agents as Optimizable Graphs. In *International Conference on Machine Learning*. PMLR, 62743–62767.

## A Algorithm

The two-stage offline training and online deployment of TACITFLOW are described in Algorithm 1.

## B Workflow Collection and Knowledge Graph Construction

To support knowledge-guided machine learning engineering, we build a workflow knowledge graph from expert competition solutions across multiple data modalities. Our objective is to convert raw notebook solutions into reusable structured workflows that preserve both the semantics of individual steps and the dependency patterns among them. Starting from public competition notebooks, we extract workflow traces, canonicalize heterogeneous step descriptions, and merge them into modality-specific directed graphs. The resulting corpus spans image, tabular, and text machine learning tasks and provides a structured substrate for downstream representation learning, workflow ranking, and agentic refinement.

### B.1 Workflow Collection

For each target modality, we begin from the public Kaggle competition catalog and issue a curated list of seed keywords describing the modality task families (e.g., classification, detection, and segmentation for image; sentiment, translation, and question-answering for text; regression and recommendation for tabular). Returned competitions are filtered to those whose problem statement matches the target modality, and the union is taken across seeds. To prevent benchmark contamination, we exclude all competitions that appear in MLE-Bench before crawling kernels or constructing ranking labels. Thus, both the workflow graph and the supervision used by the workflow ranking model are built from competitions disjoint from

**Algorithm 1** TACITFLOW: Knowledge-Grounded Machine Learning Engineering**Initialization:** Task  $t$ , dev data  $\mathcal{D}_{\text{dev}}$ , workflow graph  $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ , budget  $B$ **Input:** Best accepted solution  $\hat{c}$ 

```

1: ▷ Stage 1a: Contrastive node encoder pretraining
2: Generate augmented views  $\tilde{\mathcal{G}}_1, \tilde{\mathcal{G}}_2$  via edge dropping and feature masking
3: Train encoder  $f_\theta$  on  $(\tilde{\mathcal{G}}_1, \tilde{\mathcal{G}}_2)$  with the contrastive objective in (4.1)
4: Compute pretrained node embeddings  $\{\mathbf{h}_v\}_{v \in \mathcal{V}}$ 
5: ▷ Stage 1b: Task-conditioned workflow ranking
6: Form per-competition workflow rankings  $\{\sigma_c\}_{c \in \mathcal{C}}$  from leaderboards
7: Train ranker  $s_\phi$  on  $\{\sigma_c\}$  via the ListMLE objective in (4.2)
8: ▷ Stage 2: Online knowledge-grounded refinement
9: Initialize baseline  $\hat{c} \leftarrow c_0$ ,  $\hat{m} \leftarrow -\infty$ , history  $\mathcal{H} \leftarrow \emptyset$ , edit memory  $\mathcal{L} \leftarrow \emptyset$ 
10: while budget  $B$  not exhausted do
11:   Extract structured workflow  $w \leftarrow \text{Extract}(\hat{c})$ 
12:   for each step  $v \in \mathcal{V}_w$  do
13:     Map  $v$  to embedding space via exact / fuzzy / constructed matching
14:      $\mathcal{N}_{\text{ref}}(v) \leftarrow \{u \in \mathcal{V} : \text{sim}(\mathbf{h}_v, \mathbf{h}_u) \geq \delta, u \neq v\}$ 
15:   end for
16:   Construct candidates  $\tilde{\mathcal{W}}$  by applying each retrieved edit to  $w$ 
17:    $\mathcal{W}^* \leftarrow \bigcup_c \text{top-}k\{s_\phi(\tilde{w}, t) : \text{cat}(\tilde{w}) = c\}$ 
18:    $\gamma \leftarrow \text{Reason}(\mathcal{W}^*, \mathcal{L}, t)$  ▷ Reason over edits using edit memory
19:    $i \leftarrow \text{Propose}(\kappa, \mathcal{H})$ , where  $\kappa \triangleq (\mathcal{W}^*, \gamma)$  ▷ Generate, score, select/ensemble
20:    $c \leftarrow \text{Implement}(i; \hat{c})$ 
21:    $m \leftarrow \text{Evaluate}(c, \mathcal{D}_{\text{dev}})$ 
22:   if  $m > \hat{m}$  then
23:      $\hat{c} \leftarrow c$ ,  $\hat{m} \leftarrow m$ 
24:   end if
25:    $\mathcal{H} \leftarrow \mathcal{H} \cup \{(i, c, m)\}$ 
26:    $\ell \leftarrow \text{Summarize}(\tilde{\mathcal{W}}, i, \Delta c, m)$  ▷ Proposal, code change, validation outcome
27:    $\mathcal{L} \leftarrow \mathcal{L} \cup \{\ell\}$ 
28: end while
29: return  $\hat{c}$ 

```

the evaluation tasks. This ensures that TACITFLOW learns transferable workflow preferences rather than task-specific solutions or leaderboard information from the benchmark competitions.

For every retained competition, we crawl the public kernels, tag each one as a *solution*, *scored*, or *unscored* kernel from public metadata, and concatenate the kernels in that order into a per-modality link list. The resulting position of a workflow within its own source competition acts as a relative performance label for the downstream ranking model.

We collect workflows from public competition solutions across three modalities: image, tabular, and text. In total, the corpus contains 131 competitions and 2,681 workflows, comprising 883 image workflows from 48 competitions, 1,226 tabular workflows from 55 competitions, and 572 text workflows from 28 competitions. The competitions are disjoint across modalities, which allows each graph to capture modality-specific workflow regularities without overlap in source tasks. Across all modalities, the collected workflows exhibit substantial scale and structural diversity. In total, the corpus contains 10,126 workflow steps and 32,920 directed edges. The average workflow length is relatively consistent across modalities, ranging from 18.3 to 20.4 steps, while the average number of workflows per competition ranges from 18.4 to 22.3.

The workflow corpus covers a diverse range of machine learning settings, including classification, regression, detection, segmentation, retrieval, generation, forecasting, clustering, and ranking. It also spans multiple application domains such as healthcare, earth science, automotive, commerce, finance, education, and information retrieval. This diversity is important for learning reusable workflow components that generalize across tasks while preserving modality-specific structure. The detailed list is provided in Table 6.

## B.2 Workflow Extraction and Canonicalization

Notebook solutions often contain noisy execution traces, helper cells, debugging code, and repeated exploratory steps. To obtain reusable workflow units, we canonicalize each workflow into a sequence of semantically meaningful pipeline steps. Each step is represented as a structured node with five attributes: category: coarse pipeline stage, technique: main operation or method name, code: associated code snippet, notes: natural language description. To standardize workflows across notebooks, we use six coarse stage categories: Data Read, Data Preparing, Modeling, Training, Evaluation, and Analysis. This shared vocabulary reduces superficial variation across notebooks while preserving the functional role.

**Table 6: Kaggle competitions used for workflow collection.**

| Modality        | Competition IDs                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|-----------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Image<br>(48)   | airbus-ship-detection; benetech-making-graphs-accessible; carvana-image-masking-challenge; cifar-10; cvpr-2018-autonomous-driving; data-science-bowl-2017; data-science-bowl-2018; diabetic-retinopathy-classification-2; diabetic-retinopathy-detection; digit-recognizer; dstl-satellite-imagery-feature-detection; generative-dog-images; global-wheat-detection; google-ai-open-images-object-detection-track; google-ai-open-images-visual-relationship-track; google-universal-image-embedding; hpa-single-cell-image-classification; human-protein-atlas-image-classification; image-matching-challenge-2022; imaterialist-challenge-fashion-2018; imaterialist-challenge-furniture-2018; imet-2019-fgvc6; inclusive-images-challenge; isic-2024-challenge; jaguar-re-id; landmark-recognition-2019; landmark-recognition-challenge; landmark-retrieval-2020; landmark-retrieval-2021; landmark-retrieval-challenge; mayo-clinic-strip-ai; open-images-2019-object-detection; open-images-2019-visual-relationship; petfinder-adoption-prediction; pku-autonomous-driving; planet-understanding-the-amazon-from-space; planttraits2024; recodai-luc-scientific-image-forgery-detection; recursion-cellular-image-classification; rsna-intracranial-aneurysm-detection; sartorius-cell-instance-segmentation; sp-society-camera-model-identification; state-farm-distracted-driver-detection; tensorflow-great-barrier-reef; tpu-getting-started; ultrasound-nerve-segmentation; understanding_cloud_organization; vesuvius-challenge-surface-detection                                                                                                                                 |
| Text (28)       | asap-aes; avito-duplicate-ads-detection; classification-of-math-problems-by-kasut-academy; coleridgeinitiative-show-us-the-data; commonlit-evaluate-student-summaries; commonlitreadabilityprize; crowdfunder-search-relevance; deep-past-initiative-machine-translation; eedi-mining-misconceptions-in-mathematics; fake-or-real-the-impostor-hunt; feedback-prize-2021; feedback-prize-effectiveness; feedback-prize-english-language-learning; gendered-pronoun-resolution; jigsaw-toxic-severity-rating; llm-agentic-legal-information-retrieval; llm-classification-finetuning; llm-detect-ai-generated-text; llm-prompt-recovery; llms-you-cant-please-them-all; map-charting-student-math-misunderstandings; nbme-score-clinical-patient-notes; nlp-getting-started; quora-question-pairs; sentiment-analysis-on-movie-reviews; super-ai-ss-5-named-entity-recognition; tradeshift-text-classification; uspto-explainable-ai                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |
| Tabular<br>(55) | car-becho-paisa-paao; cat-in-the-dat; cat-in-the-dat-ii; equity-post-HCT-survival-predictions; home-depot-product-search-relevance; house-prices-advanced-regression-techniques; hull-tactical-market-prediction; instacart-market-basket-analysis; linking-writing-processes-to-writing-quality; mlsp-2014-mri; optiver-realized-volatility-prediction; otto-group-product-classification-challenge; otto-recommender-system; playground-series-s3e1; playground-series-s3e11; playground-series-s3e12; playground-series-s3e13; playground-series-s3e17; playground-series-s3e21; playground-series-s3e26; playground-series-s3e3; playground-series-s3e5; playground-series-s3e6; playground-series-s3e7; playground-series-s3e8; playground-series-s4e2; playground-series-s4e3; playground-series-s4e6; playground-series-s5e11; playground-series-s5e2; playground-series-s5e3; playground-series-s5e4; playground-series-s5e5; playground-series-s5e6; playground-series-s5e7; playground-series-s5e8; playground-series-s6e1; stock-pledge-financing-default-prediction-2026; tabular-playground-series-apr-2022; tabular-playground-series-aug-2021; tabular-playground-series-aug-2022; tabular-playground-series-feb-2021; tabular-playground-series-feb-2022; tabular-playground-series-jan-2021; tabular-playground-series-jan-2022; tabular-playground-series-jul-2022; tabular-playground-series-jun-2021; tabular-playground-series-jun-2022; tabular-playground-series-mar-2022; tabular-playground-series-may-2021; tabular-playground-series-nov-2022; tabular-playground-series-oct-2022; tabular-playground-series-sep-2021; tabular-playground-series-sep-2022; titanic |

Extraction proceeds in two passes. In the first pass, each kernel is parsed into an ordered sequence of code cells, and a single LLM call converts these cells into the structured workflow described above. The prompt asks the LLM to describe the actual operation performed by each cell. If a cell contains multiple distinct operations, it is split into multiple workflow nodes. In the second pass, we canonicalize technique names across notebooks. Variants that refer to the same method are grouped into a single canonical key using a lightweight LLM-based canonicalization step. We save these canonicalization results and reuse them in later runs, so the same technique variants are mapped consistently. The resulting nodes are also well populated: most nodes contain a code snippet, a natural language note, and a canonical technique name. Thus, each node combines categorical, textual, and code-based information, making it naturally multi-modal.

### B.3 Knowledge Graph Construction

Given the canonicalized workflows, we construct a directed workflow knowledge graph for each modality. For a modality-specific graph  $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ , each node  $v \in \mathcal{V}$  corresponds to a canonical workflow step, and each directed edge  $(u, v) \in \mathcal{E}$  indicates that step  $v$  follows step  $u$  in at least one workflow. Thus, edges encode stepwise dependencies and execution order as observed in

expert solutions. A workflow  $w$  is represented as a directed subgraph  $\mathcal{G}_w = (\mathcal{V}_w, \mathcal{E}_w) \subseteq \mathcal{G}$ , where  $\mathcal{V}_w$  is the set of workflow steps and  $\mathcal{E}_w$  contains the observed transitions among them. Beyond intra-workflow transitions, we align steps that implement the same canonical technique across notebooks and merge them into shared graph nodes. The resulting graph aggregates cross-workflow dependencies and reveals reusable pipeline patterns beyond individual notebooks. The structural composition of the graph reveals that preprocessing dominates all three modalities, followed by modeling, evaluation, and training. This is consistent with practical competition pipelines, where data preparation often accounts for a large fraction of solution effort.

### B.4 Graph Cleaning and Quality Control

Notebook-derived workflows over-produce nodes: surface-form drift in technique names, near-synonyms, and competition-specific plumbing all leave nodes that are technically distinct but functionally identical or non-transferable. The code field, which is consumed by a mean-pooled code encoder for retrieval, is also sensitive to boilerplate noise. We therefore apply a sequence of cleaning passes before finalizing the corpus.

**Embedding-based deduplication.** We compute sentence embeddings of the concatenated technique name and notes for each node,

**Table 7: Step category distribution across modalities. Values are counts with percentages in parentheses.**

| Category       | Image         | Tabular       | Text          |
|----------------|---------------|---------------|---------------|
| Data Preparing | 1,635 (46.2%) | 1,950 (50.8%) | 1,240 (45.0%) |
| Modeling       | 936 (26.5%)   | 781 (20.4%)   | 785 (28.5%)   |
| Evaluation     | 565 (16.0%)   | 549 (14.3%)   | 361 (13.1%)   |
| Training       | 388 (11.0%)   | 488 (12.7%)   | 353 (12.8%)   |
| Analysis       | 11 (0.3%)     | 67 (1.7%)     | 14 (0.5%)     |

and merge node pairs whose cosine similarity exceeds a conservative threshold. This pass removes near-duplicates that survive canonicalization, e.g., minor lexical variants of learning-rate schedules or splitting strategies.

**Singleton classification.** Nodes supported by only a single notebook are the dominant source of non-transferable noise. We classify each such node with multiple LLM votes under complementary prompt variants and apply a consensus rule, with disagreements surfaced for manual review. The classifier labels each node as keep (a generic technique likely to recur, e.g., a standard augmentation or loss), drop (competition-specific plumbing, infrastructure code, or a verbose variant already covered by an existing multi-use node), or merge (folded into a named target node). For the image modality, deduplication and singleton filtering together reduce the graph from 5,475 raw nodes to 3,536 deployed nodes, while preserving all 48 competitions and all 883 workflows; analogous reductions are obtained for text and tabular.

**Code attachment and quality control.** After constructing the canonical workflow graph, we attach a representative code snippet to each surviving node by revisiting its source notebooks. Each occurrence of a node contributes one candidate snippet. For nodes with multiple occurrences, an LLM selects the candidate that best represents the node’s canonical technique and functional role. As the retrieval module embeds the code field of each node, the attached snippet should contain technique-specific signal. To ensure this, we replace full-node rewriting with a lightweight three-stage quality-control pipeline. In **SELECT**, for nodes with multiple source occurrences, an LLM selects the candidate snippet that best represents the node’s canonical technique and functional role, favoring concise snippets that are already informative. In **GENERATE**, for nodes whose observed candidates are missing, empty, or trivially uninformative, an LLM generates a short Python snippet that directly expresses the intended operation. The prompt discourages imports, wrappers, and unrelated setup code. In **REVIEW**, for shared nodes, an LLM makes a binary **KEEP/REWRITE** decision. The review prompt favors keeping the original snippet when it contains sufficient technique signal and rewrites only when the snippet is too noisy or uninformative.

Overall, our construction process is designed to preserve three properties that are critical for downstream agent use: (i) *semantic richness*, by retaining technique names, code snippets, and natural language notes at the node level; (ii) *structural faithfulness*, by encoding execution order as directed dependencies; and (iii) *reusability*, by canonicalizing step names and merging repeated solution patterns across workflows. The resulting knowledge graph provides a compact and structured view of expert machine learning workflows. Its nodes capture reusable building blocks such as

**Table 8: Mapping between shortened competition names used in the main text and official MLE-Bench competition IDs.**

| Short name         | MLE-Bench competition ID                      |
|--------------------|-----------------------------------------------|
| dog-breed          | dog-breed-identification                      |
| plant-pathology-20 | plant-pathology-2020-fgvc7                    |
| dogs-vs-cats       | dogs-vs-cats-redux-kernels-edition            |
| tgs-salt           | tgs-salt-identification-challenge             |
| uw-gi-tract        | uw-madison-gi-tract-image-segmentation        |
| kuzushiji          | kuzushiji-recognition                         |
| denoising-dirty    | denoising-dirty-documents                     |
| spooky-author      | spooky-author-identification                  |
| detect-insults     | detecting-insults-in-social-commentary        |
| jigsaw-toxic       | jigsaw-toxic-comment-classification-challenge |
| us-patent          | us-patent-phrase-to-phrase-matching           |
| essay-score-2      | learning-agency-lab-automated-essay-scoring-2 |
| stanford-vaccine   | stanford-covid-vaccine                        |
| text-norm-russian  | text-normalization-challenge-russian-language |
| random-acts        | random-acts-of-pizza                          |
| tps-dec-2021       | tabular-playground-series-dec-2021            |
| nomad2018          | nomad2018-predict-transparent-conductors      |
| tps-may-2022       | tabular-playground-series-may-2022            |

preprocessing operations, model classes, training strategies, and evaluation routines, while its edges capture how these blocks are composed into valid pipelines. Full workflow traces are also preserved, allowing solutions to be retrieved, ranked, and adapted for new tasks. This structure is especially useful for machine learning engineering agents. Instead of exploring the solution space from scratch, an agent can reason over existing workflow components and their relations, retrieve promising pipeline structures, and modify them transparently. In this sense, the workflow graph serves not only as a memory of past expert solutions, but also as a structured search space for knowledge-guided refinement.

## B.5 Ablation of Offline Training Components

In this subsection, we provide additional ablation studies regarding different components as well as feature modalities to assess the downstream ranking performance. Table 12 reports the full results across five metrics for in-corpus data and Table 13 reports those for out-of-corpus. We also provide other ablations including the feature importance (technique name, code, notes) across three image, text, and tabular modalities, in Table 9, 10, 11 of Appendix B.5.

## C Prompts

We first provide the key components of prompts for online MLE tasks, including edit reasoning in Figure 5, proposal agent with edit guidance in Figure 6, and the coding agent in Figure 7. Second, we provide the necessary prompts for workflow extraction from Kaggle solutions and workflow knowledge graph building in Figure 8 and Figure 9. Lastly, we provide the prompts for code feature cleansing in Figure 10.

## D Full Demonstration

Please refer to this link for the full demonstration on *plant-pathology-2020-fgvc7*: <https://anonymous.4open.science/r/TacitFlow-2A19/README.md>

**Table 9: Image modality: ablation of node-level input features (in-corpus, 5-fold CV  $\times$  3 seeds).**

| Features |       |       |      | Metrics                             |                                     |                                     |                                     |                                     |
|----------|-------|-------|------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|
| Cat.     | Tech. | Notes | Code | NDCG@5                              | NDCG@10                             | MRR                                 | Top-1                               | Top-3                               |
| ✓        | ✓     | ✗     | ✗    | $0.872 \pm 0.032$                   | $0.884 \pm 0.024$                   | $0.440 \pm 0.138$                   | $0.236 \pm 0.162$                   | $0.532 \pm 0.176$                   |
| ✓        | ✗     | ✓     | ✗    | $0.874 \pm 0.035$                   | $0.889 \pm 0.025$                   | $0.435 \pm 0.127$                   | $0.230 \pm 0.163$                   | $0.530 \pm 0.160$                   |
| ✓        | ✗     | ✗     | ✓    | $0.878 \pm 0.035$                   | $0.888 \pm 0.029$                   | <u><math>0.474 \pm 0.127</math></u> | <u><math>0.272 \pm 0.152</math></u> | $0.599 \pm 0.157$                   |
| ✓        | ✓     | ✓     | ✗    | $0.882 \pm 0.041$                   | <u><math>0.896 \pm 0.027</math></u> | $0.469 \pm 0.129$                   | $0.259 \pm 0.149$                   | $0.599 \pm 0.182$                   |
| ✓        | ✓     | ✗     | ✓    | $0.884 \pm 0.033$                   | $0.891 \pm 0.024$                   | $0.472 \pm 0.128$                   | $0.238 \pm 0.181$                   | <b><math>0.664 \pm 0.145</math></b> |
| ✓        | ✗     | ✓     | ✓    | <u><math>0.885 \pm 0.029</math></u> | $0.894 \pm 0.022$                   | $0.451 \pm 0.124$                   | $0.246 \pm 0.159$                   | $0.555 \pm 0.170$                   |
| ✓        | ✓     | ✓     | ✓    | <b><math>0.897 \pm 0.037</math></b> | <b><math>0.904 \pm 0.022</math></b> | <b><math>0.488 \pm 0.132</math></b> | <b><math>0.300 \pm 0.150</math></b> | <u><math>0.612 \pm 0.200</math></u> |

**Table 10: Text modality: ablation of node-level input features (in-corpus, 5-fold CV  $\times$  3 seeds).**

| Features |       |       |      | Metrics                             |                                     |                                     |                                     |                                     |
|----------|-------|-------|------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|
| Cat.     | Tech. | Notes | Code | NDCG@5                              | NDCG@10                             | MRR                                 | Top-1                               | Top-3                               |
| ✓        | ✓     | ✗     | ✗    | $0.791 \pm 0.066$                   | $0.777 \pm 0.047$                   | $0.385 \pm 0.202$                   | $0.185 \pm 0.244$                   | $0.378 \pm 0.190$                   |
| ✓        | ✗     | ✓     | ✗    | $0.793 \pm 0.061$                   | $0.780 \pm 0.045$                   | $0.359 \pm 0.121$                   | $0.178 \pm 0.155$                   | $0.358 \pm 0.143$                   |
| ✓        | ✗     | ✗     | ✓    | $0.797 \pm 0.061$                   | $0.778 \pm 0.052$                   | $0.392 \pm 0.165$                   | $0.193 \pm 0.200$                   | $0.453 \pm 0.233$                   |
| ✓        | ✓     | ✓     | ✗    | <u><math>0.804 \pm 0.067</math></u> | $0.778 \pm 0.057$                   | <u><math>0.407 \pm 0.174</math></u> | <u><math>0.213 \pm 0.205</math></u> | $0.469 \pm 0.232$                   |
| ✓        | ✓     | ✗     | ✓    | $0.791 \pm 0.068$                   | <u><math>0.783 \pm 0.051</math></u> | $0.360 \pm 0.099$                   | $0.176 \pm 0.106$                   | $0.391 \pm 0.132$                   |
| ✓        | ✗     | ✓     | ✓    | $0.783 \pm 0.066$                   | $0.774 \pm 0.046$                   | $0.394 \pm 0.095$                   | $0.200 \pm 0.133$                   | <u><math>0.473 \pm 0.135</math></u> |
| ✓        | ✓     | ✓     | ✓    | <b><math>0.833 \pm 0.051</math></b> | <b><math>0.813 \pm 0.039</math></b> | <b><math>0.434 \pm 0.139</math></b> | <b><math>0.216 \pm 0.165</math></b> | <b><math>0.529 \pm 0.258</math></b> |

**Table 11: Tabular modality: ablation of node-level input features (in-corpus, 5-fold CV  $\times$  3 seeds).**

| Features |       |       |      | Metrics                             |                                     |                                     |                                     |                                     |
|----------|-------|-------|------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|
| Cat.     | Tech. | Notes | Code | NDCG@5                              | NDCG@10                             | MRR                                 | Top-1                               | Top-3                               |
| ✓        | ✓     | ✗     | ✗    | $0.768 \pm 0.066$                   | $0.764 \pm 0.051$                   | $0.466 \pm 0.115$                   | $0.279 \pm 0.112$                   | $0.588 \pm 0.193$                   |
| ✓        | ✗     | ✓     | ✗    | $0.752 \pm 0.100$                   | $0.752 \pm 0.071$                   | $0.443 \pm 0.184$                   | $0.236 \pm 0.159$                   | $0.594 \pm 0.318$                   |
| ✓        | ✗     | ✗     | ✓    | $0.801 \pm 0.043$                   | <u><math>0.785 \pm 0.036</math></u> | $0.517 \pm 0.102$                   | <u><math>0.333 \pm 0.127</math></u> | $0.636 \pm 0.156$                   |
| ✓        | ✓     | ✓     | ✗    | $0.782 \pm 0.080$                   | $0.773 \pm 0.051$                   | $0.494 \pm 0.164$                   | $0.303 \pm 0.171$                   | $0.642 \pm 0.254$                   |
| ✓        | ✓     | ✗     | ✓    | <u><math>0.815 \pm 0.054</math></u> | $0.784 \pm 0.038$                   | <b><math>0.543 \pm 0.090</math></b> | <u><math>0.333 \pm 0.108</math></u> | <u><math>0.703 \pm 0.157</math></u> |
| ✓        | ✗     | ✓     | ✓    | $0.784 \pm 0.081$                   | $0.773 \pm 0.054$                   | $0.459 \pm 0.140$                   | $0.242 \pm 0.131$                   | $0.642 \pm 0.231$                   |
| ✓        | ✓     | ✓     | ✓    | <b><math>0.814 \pm 0.058</math></b> | <b><math>0.790 \pm 0.044</math></b> | <u><math>0.540 \pm 0.109</math></u> | <b><math>0.342 \pm 0.113</math></b> | <b><math>0.720 \pm 0.215</math></b> |



**Table 12: Model ablation study of TacitFlow (in-corpus, 5-fold CV  $\times$  3 seeds).**

| Modality | Condition       | NDCG@5                              | NDCG@10                             | MRR                                 | Top-1                               | Top-3                               |
|----------|-----------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|
| Image    | w/o Pretraining | 0.888 $\pm$ 0.032                   | 0.895 $\pm$ 0.021                   | 0.449 $\pm$ 0.133                   | 0.228 $\pm$ 0.149                   | 0.584 $\pm$ 0.195                   |
|          | w/o VICReg      | 0.880 $\pm$ 0.026                   | 0.889 $\pm$ 0.018                   | 0.450 $\pm$ 0.118                   | 0.242 $\pm$ 0.155                   | 0.574 $\pm$ 0.128                   |
|          | w/o Attention   | 0.889 $\pm$ 0.022                   | 0.894 $\pm$ 0.020                   | 0.459 $\pm$ 0.122                   | 0.241 $\pm$ 0.153                   | 0.575 $\pm$ 0.159                   |
|          | Ours            | <b>0.897 <math>\pm</math> 0.037</b> | <b>0.904 <math>\pm</math> 0.022</b> | <b>0.488 <math>\pm</math> 0.132</b> | <b>0.300 <math>\pm</math> 0.150</b> | <b>0.612 <math>\pm</math> 0.180</b> |
| Text     | w/o Pretraining | 0.786 $\pm$ 0.061                   | 0.776 $\pm$ 0.051                   | 0.315 $\pm$ 0.110                   | 0.125 $\pm$ 0.135                   | 0.324 $\pm$ 0.181                   |
|          | w/o VICReg      | 0.804 $\pm$ 0.067                   | 0.782 $\pm$ 0.054                   | 0.427 $\pm$ 0.073                   | 0.203 $\pm$ 0.141                   | 0.522 $\pm$ 0.135                   |
|          | w/o Attention   | 0.826 $\pm$ 0.060                   | 0.800 $\pm$ 0.046                   | 0.407 $\pm$ 0.104                   | 0.189 $\pm$ 0.142                   | <b>0.547 <math>\pm</math> 0.181</b> |
|          | Ours            | <b>0.833 <math>\pm</math> 0.051</b> | <b>0.813 <math>\pm</math> 0.039</b> | <b>0.434 <math>\pm</math> 0.139</b> | <b>0.216 <math>\pm</math> 0.165</b> | 0.529 $\pm$ 0.258                   |
| Tabular  | w/o Pretraining | 0.775 $\pm$ 0.093                   | 0.762 $\pm$ 0.059                   | 0.493 $\pm$ 0.150                   | 0.280 $\pm$ 0.141                   | 0.673 $\pm$ 0.231                   |
|          | w/o VICReg      | 0.795 $\pm$ 0.055                   | 0.777 $\pm$ 0.040                   | 0.534 $\pm$ 0.113                   | 0.333 $\pm$ 0.143                   | 0.709 $\pm$ 0.204                   |
|          | w/o Attention   | 0.769 $\pm$ 0.079                   | 0.763 $\pm$ 0.050                   | 0.475 $\pm$ 0.166                   | 0.273 $\pm$ 0.159                   | 0.636 $\pm$ 0.220                   |
|          | Ours            | <b>0.814 <math>\pm</math> 0.058</b> | <b>0.790 <math>\pm</math> 0.044</b> | <b>0.540 <math>\pm</math> 0.109</b> | <b>0.342 <math>\pm</math> 0.113</b> | <b>0.720 <math>\pm</math> 0.215</b> |

**Table 13: Model ablation study of TacitFlow (OOD Evaluation on MLE-Bench runs, 3 seeds).**

| Modality | Condition       | NDCG@5                              | NDCG@10                             | MRR                                 | Top-1                               | Top-3                               |
|----------|-----------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|-------------------------------------|
| Image    | w/o Pretraining | 0.645 $\pm$ 0.019                   | 0.721 $\pm$ 0.006                   | 0.258 $\pm$ 0.012                   | 0.060 $\pm$ 0.045                   | 0.202 $\pm$ 0.110                   |
|          | w/o VICReg      | 0.650 $\pm$ 0.002                   | 0.736 $\pm$ 0.007                   | 0.256 $\pm$ 0.008                   | 0.073 $\pm$ 0.015                   | 0.260 $\pm$ 0.015                   |
|          | w/o Attention   | 0.675 $\pm$ 0.022                   | <b>0.763 <math>\pm</math> 0.006</b> | 0.242 $\pm$ 0.008                   | 0.073 $\pm$ 0.015                   | 0.271 $\pm$ 0.059                   |
|          | Ours            | <b>0.688 <math>\pm</math> 0.016</b> | <b>0.763 <math>\pm</math> 0.016</b> | <b>0.284 <math>\pm</math> 0.025</b> | <b>0.183 <math>\pm</math> 0.047</b> | <b>0.344 <math>\pm</math> 0.026</b> |
| Text     | w/o Pretraining | 0.767 $\pm$ 0.021                   | 0.825 $\pm$ 0.009                   | 0.315 $\pm$ 0.025                   | 0.079 $\pm$ 0.022                   | 0.508 $\pm$ 0.011                   |
|          | w/o VICReg      | 0.746 $\pm$ 0.025                   | 0.791 $\pm$ 0.014                   | 0.391 $\pm$ 0.020                   | 0.230 $\pm$ 0.030                   | 0.413 $\pm$ 0.040                   |
|          | w/o Attention   | 0.724 $\pm$ 0.017                   | 0.794 $\pm$ 0.008                   | 0.389 $\pm$ 0.028                   | 0.246 $\pm$ 0.049                   | 0.421 $\pm$ 0.045                   |
|          | Ours            | <b>0.783 <math>\pm</math> 0.023</b> | <b>0.829 <math>\pm</math> 0.019</b> | <b>0.519 <math>\pm</math> 0.008</b> | <b>0.381 <math>\pm</math> 0.019</b> | <b>0.587 <math>\pm</math> 0.092</b> |
| Tabular  | w/o Pretraining | 0.561 $\pm$ 0.013                   | 0.703 $\pm$ 0.011                   | 0.178 $\pm$ 0.020                   | 0.000 $\pm$ 0.000                   | 0.083 $\pm$ 0.059                   |
|          | w/o VICReg      | 0.567 $\pm$ 0.042                   | 0.720 $\pm$ 0.020                   | 0.194 $\pm$ 0.012                   | 0.000 $\pm$ 0.000                   | 0.208 $\pm$ 0.059                   |
|          | w/o Attention   | 0.638 $\pm$ 0.019                   | 0.788 $\pm$ 0.018                   | 0.281 $\pm$ 0.027                   | 0.083 $\pm$ 0.059                   | <b>0.375 <math>\pm</math> 0.000</b> |
|          | Ours            | <b>0.650 <math>\pm</math> 0.024</b> | <b>0.791 <math>\pm</math> 0.027</b> | <b>0.433 <math>\pm</math> 0.023</b> | <b>0.333 <math>\pm</math> 0.059</b> | <b>0.375 <math>\pm</math> 0.000</b> |

| Key Prompt of Edit Reasoning                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>**System.**</b></p> <p>You are an experienced ML research advisor reviewing a candidate-edit menu for a Kaggle [MODALITY] pipeline. The menu was retrieved by an external workflow ranker from a knowledge base of past Kaggle solutions, then will be passed to a downstream hypothesis generator.</p> <p>Your role: ADVISORY, not prescriptive. The downstream generator decides freely. Your job is to enrich the menu with one short advisory note per option so the generator can engage more thoughtfully.</p> <p>Allowed actions per option:</p> <p>(a) Caveat — flag if the option resembles a recently failed attempt or was rejected in an earlier loop of this run.</p> <p>(b) Pair — when two options would naturally combine given the base workflow, mention this briefly.</p> <p>(c) Frame — provide context on when this technique typically helps for this scenario (task type, metric, data size, time budget).</p> <p>Forbidden:</p> <ul style="list-style-type: none"> <li>- Do NOT tell the generator to adopt or reject a specific option.</li> <li>- Do NOT rewrite option content, rescore, or rerank.</li> <li>- Do NOT invent options outside the menu.</li> <li>- Do NOT generate a hypothesis.</li> </ul> | <p>Output: reproduce the menu exactly (preserve category headers and option numbering). Under each option add one short advisory paragraph (2-3 sentences). Omit the paragraph if you have nothing useful to say. Keep total output under 800 words.</p> <p><b>**User.**</b></p> <p><b># Scenario Summary</b><br/>[SCENARIO_SUMMARY]</p> <p><b># Base Workflow</b><br/>[BASE_WORKFLOW]</p> <p><b># Last Loop Outcome</b><br/>[LAST_OUTCOME]</p> <p><b># Recent Lessons from this Run (most recent first)</b><br/>[MEMORY_RECENT]</p> <p><b># Candidate Menu (ranker output, ranked by category)</b><br/>[CANDIDATES]</p> <p>Reproduce the menu and add a short Advisory paragraph below each option where useful. Stay strictly advisory.</p> |

**Figure 5: Prompt used by the proposal agent to reason over ranked workflow edits.**

| Key Prompt of Proposal Agent with Edit Guidance                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>**System.**</b><br/>You are a Kaggle Grandmaster and expert ML engineer with deep expertise in statistics, machine learning, and competition optimization. The user is iteratively improving a Kaggle competition implementation. Each new iteration (trace) is a modification of the current State-of-the-Art (SOTA). If a new trace surpasses the current SOTA, it becomes the new SOTA. Otherwise, it's a failed experiment.</p> <p>You will be provided with:</p> <ol style="list-style-type: none"> <li>1. A detailed competition description and identified challenges from history.</li> <li>2. A history of previous successful and failed experiments and their feedback.</li> <li>3. The current SOTA implementation and feedback.</li> <li>4. A set of Ranked Candidate Edits derived from a knowledge base of Kaggle workflows. Each edit is a technique-level replacement in the current workflow, scored by a trained workflow quality model.</li> </ol> <p>Task 1 - Hypothesis Proposal: for each Identified Challenge, propose one specific, testable hypothesis.</p> <p>Task 2 - Hypothesis Evaluation: assign each hypothesis a component tag and score it on five 1-10 dimensions (alignment, expected impact, novelty, feasibility, risk-reward balance).</p> <p>Workflow Improvement Suggestions from Knowledge Base.<br/>You are also provided with ranked candidate edits derived from successful Kaggle workflows. These suggest technique replacements at specific workflow positions, scored by a trained workflow quality model. If a suggested edit is compatible with the current competition's data type, model framework, and pipeline constraints, consider adopting it as the basis for one of your hypotheses.</p> | <p><b>**User.**</b><br/><b># Scenario Description</b><br/>[SCENARIO_DESC]</p> <p><b># Previous Experiments and Feedbacks</b><br/>[EXP_AND_FEEDBACK_LIST]</p> <p><b># Current SOTA Implementation</b><br/>[SOTA_EXP_DESC]</p> <p><b># Identified Challenges</b><br/>[PROBLEMS]</p> <p><b># Ranked Candidate Edits</b><br/>The section below shows: (1) the current workflow baseline, and (2) suggested edits ranked by an external workflow ranker trained on successful Kaggle workflows.</p> <p>Instructions:</p> <ul style="list-style-type: none"> <li>- Review the current workflow steps to understand the baseline.</li> <li>- Choose at most ONE edit to implement as your core hypothesis change.</li> <li>- The edit suggestions are based on what worked well in similar competitions.</li> <li>- If you choose none, explicitly state why they are unsuitable.</li> </ul> <p>[RANKED_EDIT_SUGGESTIONS]</p> <p>At runtime, [RANKED_EDIT_SUGGESTIONS] renders as the current numbered workflow plus retrieved edits grouped by category, each formatted as "Option N: replace [OLD] → [NEW] at step [TARGET_IDX]".</p> |

Figure 6: Prompt used by the proposal agent to generate candidate hypotheses based on edit guidance.

| Key Prompt of Coding Agent                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>**System.**</b><br/>You are a grandmaster-level data scientist and machine learning engineer with deep expertise in statistics, mathematics, and computer science. Your task is to generate robust, debuggable, iteration-friendly code for data science pipelines, following a strict, stepwise process. You are working on sample datasets and your code will go through automated iterations. Design your code to be iteration-friendly with comprehensive print statements and clear debugging information.</p> <p><b># Task Description</b><br/>[TASK_DESC]</p> <p><b>## Runtime Environment</b><br/>[RUNTIME_ENVIRONMENT]</p> <p><b>## Package Information</b><br/>[PACKAGE_INFO]</p> <p><b>## Hyperparameters Specification</b><br/>Follow hyperparameter choices specified in the task unless they are unreasonable or incorrect.</p> <p><b># Specification your code should follow</b><br/>[SPEC]</p> <p><b># Workflow Overview</b></p> <p><b>## Data Loading</b><br/>- Load the dataset strictly from the configured input path. Do NOT load from the current directory.<br/>- Beyond data loading, avoid try-except blocks to hide or suppress errors. All errors should be properly diagnosed and fixed at source.</p> <p><b>## Exploratory Data Analysis</b><br/>Output a plain-text EDA section to stdout within the markers:<br/>=== Start of EDA part ===<br/>[EDA content: shape, dtypes, missing values, target distribution, feature analysis]<br/>=== End of EDA part ===</p> <p><b>## Debug Mode</b><br/>The code is invoked as `python main.py --debug`. In debug mode, sample 10% of training data and run minimum epochs, then print:<br/>=== Start of Debug Information ===<br/>debug_time: &lt;seconds&gt;</p> | <p>estimated_time: &lt;seconds for full run&gt;<br/>=== End of Debug Information ===<br/>In full mode, use all data and full epochs.</p> <p><b>## General Guidelines</b></p> <ol style="list-style-type: none"> <li>1. Code correctness is the top priority. Ensure your code is runnable and produces the expected output even if some task requirements are not fully met.</li> <li>2. Use print() for all output; do not use the logging module.</li> <li>3. Avoid all hard-coded values (e.g., fixed dataset sizes). Always use proportions for data splitting, never absolute numbers.</li> <li>4. Add informative print statements at key steps.</li> <li>5. For model training, use reasonable epoch numbers and implement early stopping. Save best model checkpoints based on validation performance.</li> <li>6. Except in debug mode, ALWAYS use all available data; do not sample or subset due to resource limitations.</li> <li>7. Do not use tqdm or progress bars.</li> <li>8. Try-except blocks are ONLY allowed when reading files. Assert statements are PROHIBITED throughout.</li> <li>9. ALWAYS use the best saved model for predictions. NEVER create dummy submissions (all 1s, random values). If training fails, report failure honestly rather than generating fake submission files.</li> <li>10. Generate the complete code, not partial code.</li> <li>11. If the task contains user instructions, follow them - they have the highest priority.</li> </ol> <p><b>**User.**</b><br/><b># Competition Information</b><br/>[COMPETITION_INFO]</p> <p><b># Data Folder Description</b><br/>[FOLDER_SPEC]</p> <p><b># Former code (if iterating)</b><br/>[LATEST_CODE]</p> <p><b>## Feedback to former code</b><br/>[LATEST_CODE_FEEDBACK]</p> <p>Before modifying the code, analyze the feedback and identify no more than three key areas requiring changes. Prioritize critical issues affecting execution, correctness, or stability; preserve existing working components.</p> |

Figure 7: Prompt used by the coding agent to generate robust, debuggable, and iterative MLE code.

| Key Prompts of Workflow Extraction and Workflow Knowledge Graph Building (Part 1)                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                     |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>**A.4 Notebook Triage.**</b><br/>You judge if the document (notebook code, GitHub README, or Kaggle solution writeup) provides enough information to understand and reproduce ONE coherent final workflow that yields the official result.</p> <p>Rules</p> <ol style="list-style-type: none"> <li>1. Suitable if the document clearly shows how data, modeling, and prediction connect to produce the final result.</li> <li>2. Ensembles, stacks, or cascades that merge into one final output count as ONE workflow.</li> <li>3. Not suitable if multiple unrelated final pipelines are presented without a single chosen one.</li> <li>4. Pure analysis or visualization without modeling/prediction -&gt; not suitable.</li> <li>5. Never infer missing steps; rely only on explicitly provided content.</li> </ol> <p>Input<br/>[NOTEBOOK_CLUES]</p> <p>Output<br/>decision: "suitable"   "not_suitable", reason: "&lt;=120 chars"</p> <p><b>**A.5 Workflow Extraction (PLAN pass).**</b><br/>You are extracting a canonical workflow plan from a Kaggle source document.</p> <p>[INSIGHTS] curated claims and techniques mined from text/discussion.</p> | <p>Derive an ordered list of pipeline steps covering the full workflow implied by [INSIGHTS]. Each step is a (category, technique) pair.</p> <p>Categories (assign exactly one per step)<br/>- Data Read   Data Preparing   Modeling   Training   Evaluation</p> <p>Coverage. Every workflow MUST include at least one step for each of the following slots, specialized strictly from [INSIGHTS]:</p> <ol style="list-style-type: none"> <li>(1) data source / datamix construction</li> <li>(2) data transformation / augmentation / normalization</li> <li>(3) feature representation</li> <li>(4) base modeling</li> <li>(5) ensemble or meta modeling (if implied)</li> <li>(6) training scheme</li> <li>(7) evaluation protocol</li> </ol> <p>Hard constraint: if [INSIGHTS] mention 3+ distinct techniques, output at least 8 pipeline steps; if 6+ items, at least 10.</p> <p>Rules</p> <ul style="list-style-type: none"> <li>- Promote any operation causing a semantic state change in the pipeline into its own step.</li> <li>- Technique names: unique, concise, lowercased, canonical.</li> <li>- Exclude generic placeholders unless a specific algorithm is named.</li> <li>- Maintain a realistic runtime ordering.</li> </ul> |

Figure 8: Prompts for workflow extraction and workflow knowledge graph construction: notebook triage, step extraction.

| Key Prompts of Workflow Extraction and Workflow Knowledge Graph Building (Part 2)                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                   |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                         |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>**A.6.1 Cluster naming.**</b><br/>You standardize step names for a cluster of similar steps.</p> <p>Goal: produce canonical (category, techniques) pairs that are precise, reproducible, and library-aligned.</p> <ul style="list-style-type: none"> <li>- If all excerpts describe one operation -&gt; output ONE canonical item.</li> <li>- If they describe distinct operations -&gt; output MULTIPLE items in logical runtime order.</li> <li>- Prefer existing canonical names. Reuse an existing name if it refers to the same technique even if naming or arguments differ. Create a new name only when no existing item captures the same library class or algorithmic intent.</li> </ul> <p>Naming: Category in {Data Read, Data Preparing, Modeling, Training, Evaluation, Analysis}. Technique = library class/method name or canonical algorithm name (sklearn.preprocessing.StandardScaler.transform, torch.optim.Adam, resnet, lightgbm). One technique per item; expand chained calls into atomic steps.</p> <p>Target cluster<br/>[CLUSTER_EXCERPTS]</p> <p>Existing relevant names<br/>[CANDIDATE_EXISTING_NAMES]</p> <p><b>**A.6.2 Action selection.**</b><br/>You decide how to handle ONE target step.</p> <ul style="list-style-type: none"> <li>- CATEGORY: [CATEGORY]</li> <li>- TECHNIQUES: [TECHNIQUES]</li> <li>- NOTE: [NOTES]</li> <li>- HYPERPARAM: [HYPHS]</li> </ul> <p>Classify into one of:</p> <ul style="list-style-type: none"> <li>- "drop" : only structural / I/O / wrapping; no semantic change to data, model, or evaluation.</li> <li>- "rename_or_split" : ambiguous, generic, user-defined, dataset-specific, or combines multiple operations.</li> <li>- "proper" : one clear, reproducible, technically meaningful process.</li> </ul> <p>When uncertain, prefer "rename_or_split" over "proper".</p> <p><b>**A.6.3 Rename/split refinement.**</b><br/>You refine ONE target step whose (category, techniques) may need rename and/or split.</p> <p>Input<br/>{ "category": "[CATEGORY]", "techniques": "[TECHNIQUES]", "notes": "[NOTES]",</p> | <p>"hyperparams": "[HYPHS]", "code": "[CODE_EXAMPLE]" }</p> <p>Rules</p> <ul style="list-style-type: none"> <li>- SPLIT when evidence indicates multiple distinct operations.</li> <li>- RENAME when the technique is generic or user-defined / dataset-specific.</li> <li>- KEEP as a single item when already specific and reproducible.</li> <li>- EXCLUDE utility/housekeeping steps with no effect on data, features, model training, or evaluation.</li> </ul> <p><b>**A.6.4 Match to global vocabulary.**</b><br/>You canonicalize ONE operation by matching a TARGET PAIR against existing canonical names, using contextual evidence.</p> <p>Target pairs (candidates for THIS operation)<br/>[CANDIDATE_PAIRS]</p> <p>Context<br/>{ "notes": "[NOTES]", "hyperparams": "[HYPHS]", "code": "[CODE_EXAMPLE]" }</p> <p>Existing relevant names (global vocabulary so far)<br/>[EXISTING_NAMES]</p> <p>Decide</p> <ol style="list-style-type: none"> <li>1) Which target pair best describes the operation.</li> <li>2) Then:       <ol style="list-style-type: none"> <li>a) REUSE an existing canonical name if it represents the same op.</li> <li>b) KEEP the chosen pair (possibly with minor normalization).</li> <li>c) CREATE a new canonical pair only if neither applies.</li> </ol> </li> <li>3) Constraint: at most one "Data Read" step per workflow.</li> </ol> <p><b>**A.7 Dataset-Task Similarity Scoring.**</b><br/>You are a meticulous dataset-task similarity judge.</p> <p>Facets (each scored 0..5, weights in parentheses)</p> <ol style="list-style-type: none"> <li>F1. Modality &amp; Channels (30%)</li> <li>F2. Unit of Analysis &amp; Domain (20%)</li> <li>F3. Supervision / Target Structure (20%)</li> <li>F4. Intrinsic Structure (spatial/temporal/hierarchy) (15%)</li> <li>F5. Distribution &amp; Data Quality (10%)</li> <li>F6. Metric &amp; Evaluation Protocol ( 5%)</li> </ol> <p>Per candidate:</p> <ol style="list-style-type: none"> <li>1) Assign facet scores f1..f6 in {0..5}.</li> <li>2) <math>S = 100 * (0.30 * f1/5 + 0.20 * f2/5 + 0.20 * f3/5 + 0.15 * f4/5 + 0.10 * f5/5 + 0.05 * f6/5)</math></li> <li>3) raw_score = round(S).</li> </ol> |

Figure 9: Prompts for workflow extraction and workflow knowledge graph construction: canonicalization and step-level similarity scoring.

| Key Prompts of Code Quality Control                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                          |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p><b>## A.8 Code-Snippet Selection</b><br/>When a knowledge-graph node aggregates multiple code snippets harvested from different notebooks, picks the single best one.</p> <p><b>**System.**</b><br/>You are a knowledge graph curator. Select the single best code snippet to represent a technique node. The snippet will be embedded by CodeT5+ (mean pooling over 512 tokens) for similarity retrieval.</p> <p>A good snippet:<br/>- Is complete, runnable Python (not a fragment or config value).<br/>- Is generic and reusable (no hardcoded competition paths or column names).<br/>- Clearly illustrates the described technique.<br/>- Starts with the most distinctive code (not imports).</p> <p>Avoid:<br/>- Very short config lines (SEED=42, n_folds=5).<br/>- Entirely commented out code.<br/>- Shell or magic commands.<br/>- Code hardcoded to one specific competition.</p> <p><b>**User.**</b><br/>Select the single best code snippet for the following node.<br/>Technique: [TECHNIQUE NAME]<br/>Category: [CATEGORY]<br/>Description: [TECHNIQUE DESCRIPTION]<br/>Candidate snippets:<br/>---- Snippet 1 ----<br/>[CANDIDATE SNIPPET 1]<br/>---- Snippet 2 ----<br/>[CANDIDATE SNIPPET 2]<br/>...<br/>Reply with ONLY the snippet number (e.g. "3"). No explanation.</p> <p><b>## A.9 Code-Snippet Generation</b><br/>When every candidate snippet for a node is unusable, synthesizes a canonical implementation from the node's metadata.</p> <p><b>**System.**</b><br/>You are a Python code generator for a machine learning knowledge graph. Given a technique description, write a short, representative Python snippet that will be embedded by CodeT5+ for retrieval.</p> <p>Rules:<br/>- Valid, runnable Python (not pseudocode).<br/>- Generic variable names (X, y, df, train, test, model).<br/>- Concise: 3 to 20 lines, focused on the core technique.<br/>- Do NOT add import statements.<br/>- Use def/class only when the technique genuinely needs a function.<br/>- Do NOT include markdown fences or explanations.<br/>- Do NOT hardcode competition-specific file paths.<br/>- Start with the most distinctive line first.</p> | <p><b>**User.**</b><br/>Generate a Python snippet for the following ML technique node.<br/>Technique: [TECHNIQUE NAME]<br/>Category: [CATEGORY]<br/>Description: [TECHNIQUE DESCRIPTION]<br/>Hyperparameters: [TECHNIQUE METADATA]<br/>Existing code: [EXISTING CODE – low quality, improve or replace]<br/>Write ONLY the Python code, no imports, no markdown fences.</p> <p><b>## A.10 Code-Snippet Review</b><br/>Passes every node's snippet through a final KEEP-or-REWRITE pass with strong KEEP bias.</p> <p><b>**System.**</b><br/>You are reviewing code for a machine learning knowledge graph. Each node represents one technique (e.g., "pca", "stratified-kfold"). The code will be embedded by CodeT5+ using mean pooling over 512 tokens for retrieval.</p> <p>Embedding-aware rules:<br/>1. Do NOT add import statements. Imports waste token budget and make all nodes look similar in embedding space.<br/>2. Do NOT wrap simple code in a function just to wrap it. A clear one-liner like 'skf = StratifiedKFold(n_splits=5, shuffle=True, random_state=42)' is MORE distinctive for retrieval than a generic wrapper. Use 'def' only when the technique genuinely needs a multi-step routine.<br/>3. Lead with the most distinctive code. The first tokens matter most for retrieval; do not open with boilerplate.<br/>4. Keep it concise. Aim for 2 to 15 lines. Every token should carry information about THIS specific technique.</p> <p>Default to KEEP. Only REWRITE when the code has a serious problem in one of these categories:<br/>- Empty, commented out, shell command, or non-Python.<br/>- Implements a DIFFERENT technique than described.<br/>- Entirely competition-specific (hardcoded paths/columns, no reusable logic).<br/>- Bloated beyond recognition (&gt;2000 chars of irrelevant code).<br/>- Just a config value with no API call (e.g., 'NFOLDS=5' alone).<br/>Write ONLY improved Python code when rewriting. No explanations, no fences.</p> <p><b>**User.**</b><br/>Technique: [TECHNIQUE NAME]<br/>Category: [CATEGORY]<br/>Description: [TECHNIQUE DESCRIPTION]<br/>Hyperparameters: [TECHNIQUE METADATA]<br/>Current code:<br/>[EXISTING CODE]<br/>Reply "KEEP" if the code acceptably illustrates the technique – even imperfect code that clearly shows the technique is fine. Reply "REWRITE" followed by improved code only if one of the problem categories applies.</p> |

Figure 10: Prompts for code-snippet quality control, including snippet selection, generation, and review.