

Task Cascades for Efficient Unstructured Data Processing

SHREYA SHANKAR, University of California Berkeley, USA

SEPANTA ZEIGHAMI, University of California Berkeley, USA

ADITYA G. PARAMESWARAN, University of California Berkeley, USA

Modern database systems allow users to query or process unstructured text or document columns using LLM-powered functions. Users can express an operation in natural language (e.g., “*identify if this review mentions billing issues*”), with the system executing the operation on each document, in a row-by-row fashion. One way to reduce cost on a batch of documents is to employ the model cascade framework: a cheap proxy model processes each document, and only uncertain cases are escalated to a more accurate, expensive oracle. However, model cascades miss important optimization opportunities; for example, often only part of a document is needed to answer a query, or other related, but simpler operations (e.g., “*is the review sentiment negative?*”, “*does the review mention money?*”) can be handled by cheap models more effectively than the original operation, while still being correlated with it.

We introduce the **task cascades** framework, which generalizes model cascades by varying not just the model, but also the *document portion* and *operation* at each stage. Our framework uses an LLM agent to generate simplified, decomposed, or otherwise related operations and selects the most relevant document portions, constructing hundreds of candidate tasks from which it assembles a task cascade. We show that optimal cascade selection is intractable via reduction from MINIMUM SUM SET COVER, but our iterative approach constructs effective cascades. We also provide an extension that offers statistical accuracy guarantees: the resulting cascade meets a user-defined accuracy target (with respect to the oracle) up to a bounded failure probability. Across eight real-world document processing tasks at a 90% target accuracy, task cascades reduce end-to-end cost by an average of 36% compared to model cascades, at a production scale.

CCS Concepts: • **Information systems** → **Data management systems**; • **Applied computing** → *Document management and text processing*; • **Computing methodologies** → *Natural language processing*.

Additional Key Words and Phrases: Unstructured data processing, semantic data processing, LLMs

ACM Reference Format:

Shreya Shankar, Sepanta Zeighami, and Aditya G. Parameswaran. 2026. Task Cascades for Efficient Unstructured Data Processing. *Proc. ACM Manag. Data* 4, 1 (SIGMOD), Article 88 (February 2026), 26 pages. <https://doi.org/10.1145/3786702>

1 Introduction

Large language models (LLMs) are increasingly being integrated as operators in data management systems, enabling users to analyze unstructured text columns more easily. This line of work spans SQL extensions to databases, providing LLM-powered *map* and *filter* operations [16, 20, 54, 65], as well as Python-based declarative frameworks for executing pipelines of LLM-powered operators [1, 40, 47, 51, 63]. In all cases, users describe an operation in natural language, with the system executing it by invoking an LLM on each text value—hereafter referred to as a *document*—in a row-by-row fashion. Consider the following example:

Authors’ Contact Information: Shreya Shankar, University of California Berkeley, USA, shreyashankar@berkeley.edu; Sepanta Zeighami, University of California Berkeley, USA, zeighami@berkeley.edu; Aditya G. Parameswaran, University of California Berkeley, USA, adityagp@berkeley.edu.



This work is licensed under a Creative Commons Attribution 4.0 International License.

© 2026 Copyright held by the owner/author(s).

ACM 2836-6573/2026/2-ART88

<https://doi.org/10.1145/3786702>

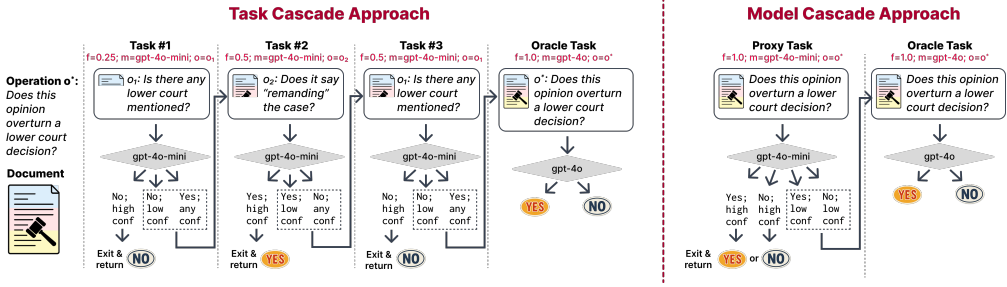


Fig. 1. Task vs. model cascade for determining whether a Supreme Court opinion overturns a lower court decision in Example 1.1. In a task cascade (left), each stage (or task) is defined by a document fraction (f), a model (m), and an operation (o). Surrogate operations may be reused at different document fractions (as in tasks 1 and 3). The final oracle task applies the original user-specified operation (o^*) on the full document ($f = 1.0$) with the oracle model ($m = gpt-4o$) if prior tasks cannot confidently resolve the input. In contrast, a model cascade (right) applies the same operation o^* to the full document at each stage, varying only the model.

Example 1.1 (Supreme Court Opinion Analysis). Suppose a legal analyst is working with a large collection of Supreme Court opinions. The analyst’s goal is to determine, for each opinion, whether it overturns a lower court decision. In this setting, each *document* is a long, complex legal text, and the analyst poses an *operation* in natural language: “Does this opinion overturn a lower court decision?”

State-of-the-art LLMs like GPT-4o and GPT-4.1 can achieve high accuracy on tasks such as Example 1.1. However, commercial LLM APIs charge by the number of input and output tokens, making large-scale analysis costly. As a result, users look for ways to cut inference costs without losing much of the accuracy provided by the best models. They are often willing to accept a modest drop in accuracy—e.g., targeting 90% of GPT-4o’s accuracy—in exchange for significant savings [47, 50]. Although cheaper models exist (such as GPT-4o-mini, which is nearly 17× cheaper than GPT-4o), they often deliver much lower accuracy.

Prior Work and Limitations. The standard approach to reduce costs of map and filter operations with machine learning models, while preserving output accuracy, is to use a *model cascade* [2, 33, 35, 36, 49]; with recent systems all adopting it as-is for LLM-powered operators [9, 47, 68]. Here, the system first applies a cheaper *proxy* model to each input document. If the proxy’s confidence in its prediction exceeds a system-selected confidence threshold, the system accepts the proxy’s output. Otherwise, the system escalates the input to a more accurate but more expensive *oracle* model (e.g., the best available LLM). Users specify a performance target (e.g., 90% of the oracle’s accuracy), and the system tunes the confidence thresholds to meet this target while minimizing cost.

However, the model cascades framework, by focusing only on swapping models, underutilizes optimization opportunities. For instance, when no cheaper model can perform the operation accurately, documents will end up being processed by the oracle. In fact, on medical or legal tasks, cheaper models may have as little as 30% of the accuracy of expensive models [17, 46, 61], with the gap being even wider for long-context reasoning [41, 66]. Moreover, model cascades always process the full document, even though many operations require only a small, relevant portion of the text.

Change the Task, Not Just the Model. Our key insight is that *effective cascades should vary the operation as well as the portion of the document*, in addition to the model. Earlier stages of the cascade can execute any operation that is easier for cheaper models—such as a prerequisite check, a decomposed sub-task, or any related operation that is correlated with the original. We refer to such

operations as *surrogate* operations: variants of the original operation that are easier for a proxy model to perform. In Example 1.1, rather than having a proxy model decide if an opinion overturns a lower court decision, we can have it check whether the opinion mentions a lower court at all. If not, we don't need to invoke the oracle.

Another complementary strategy to improve proxy model accuracy is to *prune irrelevant context* from each document before inference. For example, instead of providing the entire document to the LLM, we can supply only the sections that are semantically relevant to the operation—such as paragraphs containing “overturn,” “judgment vacated,” or their synonyms. By eliminating unrelated text, the proxy model is more likely to perform the operation both correctly *and* with lower cost [57, 67].

In this paper, we introduce the notion of a **task cascade**: an ordered sequence of *tasks*, each task defined by an operation, an LLM, and a selected fraction of the document to analyze. As illustrated in Figure 1, for Example 1.1, a task cascade may first check if there is any lower court mentioned in a short excerpt before switching over to a different surrogate, and later revisit this same operation over even more of the document—before finally applying the original operation with the oracle model on the full document if needed.

Task Cascade Challenges. However, designing an effective task cascade is challenging for several reasons. First, discovering effective surrogate operations is hard, and there is an infinitely large search space. While one can use an LLM to author a new operation, a surrogate operation is only beneficial if it is both accurately performed by a cheap model *and* if it resolves a substantial fraction of documents. A poor surrogate can potentially make the entire cascade *less* efficient than if it were omitted. Second, determining the ideal portion of the document for each task is non-trivial. The goal is to minimize token costs by providing just enough context for an accurate decision; too little text risks errors, while too much (like a full document) negates potential savings and can overwhelm cheaper models. Third, determining the best order in which to apply tasks is difficult. The number of possible sequences grows rapidly with more tasks, and each sequence can differ significantly in cost and effectiveness. A task may perform relatively well overall but poorly when applied only to the harder examples that earlier tasks did not resolve, and thus not meet accuracy constraints. Moreover, LLM APIs offer cost savings (50–90%) when multiple prompts share the same input prefix. As a result, optimizing task order is not just about accuracy and selectivity of tasks: it also involves structuring inputs to exploit reuse. ***Overall, jointly optimizing surrogate operations, document fractions, and task orders creates a complex search space that is difficult to reliably navigate.***

Effective Task Cascades. We introduce an approach for constructing effective task cascades, given a document set, user-specified operation, and a target accuracy with respect to an oracle. We address the aforementioned challenges as follows:

- **Task ordering.** Given a set of tasks, trying all orderings to find the best one quickly becomes infeasible as the number of tasks grows. We show that optimally ordering tasks is NP-HARD through a reduction from MINIMUM-SUM-SET-COVER, and inspired by approximation algorithms for this problem, we develop a greedy algorithm that sequentially adds the task that most reduces total inference cost while satisfying accuracy constraints.
- **Confidence threshold selection.** Per task in the cascade, we need to set confidence thresholds that determine whether to accept the proxy model output for that task. A naive approach sets thresholds for each task so that, by a union bound, the overall cascade meets the target accuracy with the desired probability. However, as the number of tasks increases, the naive method becomes overly conservative, making it difficult to find a feasible cascade. We develop a new threshold

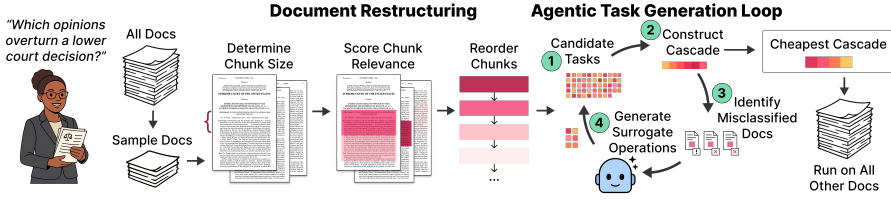


Fig. 2. Overview of our approach. A user poses a complex classification query over long documents. Our approach restructures each document to prioritize relevant chunks, then iteratively proposes surrogate operations and assembles a cost-effective task cascade to meet accuracy constraints.

adjustment algorithm, building on recent statistical results [64], which enables us to achieve accuracy guarantees with far fewer samples.

- **Surrogate operation generation.** To generate the set of tasks in the first place, a simple idea is to prompt an LLM agent (e.g., based on OpenAI’s o1) for candidate operations, but these surrogates are often not effective for smaller models, as the agent has no way to know which operations are actually easy for cheaper models to perform. Instead, we introduce an *iterative* approach, where the agent proposes surrogates, and tests and refines them based on which are actually performed well by proxies.

- **Document pruning.** Finally, per task, we want to minimize the portion of the document provided as input, while maintaining accuracy. A naive approach to pruning irrelevant portions might use embedding similarity to select document chunks most related to the user operation, similar to standard retrieval-augmented generation (RAG) techniques; however, this approach often fails for complex or long-context tasks, so we instead train a lightweight classifier—supervised by the oracle LLM—to score each chunk’s relevance and use these scores for pruning.

Overall, our contributions include:

- (1) We introduce the concept of *task cascades* for LLM-based document processing, generalizing model cascades by allowing each stage to specify a model, operation (original or surrogate), and document fraction.
- (2) We show that the problem of optimal cascade construction is NP-HARD, via a reduction from MINIMUM SUM SET COVER.
- (3) We develop an automated, agentic approach for constructing effective task cascades—jointly discovering surrogate operations, pruning documents, and greedily assembling task sequences that minimize cost while meeting accuracy targets. We show that our approach can be extended to provide guarantees.
- (4) We evaluate our method on eight complex document classification tasks drawn from Kaggle and prior LLM-based data processing research, comparing against standard model cascade baselines and analyzing the contribution of each component (task ordering, document pruning, surrogate generation). Across all workloads at a 90% target accuracy, **our approach reduces inference cost by 48.5% on average compared to model cascade baselines, and by 86.2% compared to using the oracle model alone.**

We present the problem and an overview of our approach in Section 2. Next, in Section 3, we show the hardness of optimal task ordering and describe our greedy cascade assembly algorithm and new procedure for achieving statistical accuracy guarantees. In Section 4, we describe how we prune irrelevant context in documents, followed by our agentic approach for discovering surrogate operations in Section 5. Then, in Section 6, we describe how to model the cost of task cascades. Finally, we present our evaluation in Section 7 and cover related work in Section 8.

2 Problem Setup and Approach

In this section, we formalize the notion of task cascades and give an overview of our approach to finding an efficient task cascade. A summary of all notation is provided in Table 1, for convenience.

2.1 Setup

Problem Setting. We focus on the setting where LLM-powered *map* and *filter* operations in commercial databases [16, 20, 54, 65] or recent systems [40, 47, 51] produce outputs from a fixed set of classes. Formally, the user provides a collection D of documents and a target operation o_{orig} , described in natural language. The goal is to assign each document $x \in D$ to a class c from a predefined set of classes C . The user specifies an accuracy target $\alpha \in (0, 1]$, requiring the system’s predictions to match an oracle model m_{oracle} ’s predictions on at least an α -fraction of documents.

Given a document x (or, alternatively, a fraction of the document, as described below) and operation o (either the original or a surrogate, defined below), a model m produces a score $p_m(c \mid x, o)$ for each class $c \in C$. The class c with the highest $p_m(c \mid x, o)$ is taken as the prediction of the model. For Example 1.1, x could be a court opinion, o might be “Does this opinion overturn a lower court decision?”, and $c \in \{\text{True}, \text{False}\}$. To enable early termination in a cascade, we use **confidence thresholds**: for each class $c \in C$, if the model’s confidence $p_m(c \mid x, o) \geq \tau^c$ for some threshold τ^c , the prediction for that document x is accepted; otherwise, x continues to the next stage.

Instead of always using the full document x , we may also use a portion of the document, x_f , denoting the top f fraction of x , containing the most relevant content for performing o_{orig} . Relevance is determined by a scoring function that ranks document segments by their utility for o_{orig} . So, the confidence threshold now based on $p_m(c \mid x_f, o)$. A **task** is then defined as $T_i = (m_i, o_i, f_i, \tau_i)$, where:

- m_i : model (e.g., proxy or oracle)
- o_i : operation (original or surrogate)
- f_i : fraction of the document processed
- $\tau_i = \{\tau_i^c\}_{c \in C}$: class-specific confidence thresholds

A **task cascade** is an ordered sequence of tasks $\pi = (T_1, T_2, \dots, T_k)$. At inference time, for a given document x , we evaluate each T_i in sequence: model m_i is applied to input x_{f_i} with operation o_i , yielding a predicted class c and confidence score $p_{m_i}(c \mid x_{f_i}, o_i)$. If this confidence exceeds the corresponding threshold τ_i^c , the prediction is accepted, $\text{Cascade}(\pi, x) = c$, and the cascade terminates; otherwise, processing continues to the next task. If none of the k tasks return a confident prediction, the cascade defers to the oracle task $T_{k+1} = (m_{\text{oracle}}, o_{\text{orig}}, 1, \emptyset)$, which applies o_{orig} to the full document with the oracle, with $\text{Cascade}(\pi, x)$ set to the oracle’s prediction. Thus, a document “leaves” the cascade at the first confident prediction within π , or at the oracle if all tasks defer. A traditional model cascade is thus a “restricted” task cascade where each T_i uses the same operation $o_i = o_{\text{orig}}$ and processes the entire document ($f_i = 1.0$) at each stage, varying only the model m_i .

Confidence Scores. Each model m_i defines a distribution over output classes $c \in C$, conditioned on operation o_i and document portion x_{f_i} , i.e., $p_{m_i}(c \mid x_{f_i}, o_i)$. Inputs to LLMs are tokenized; a *token* is a subword unit such as a word fragment or punctuation mark. LLMs return log-probabilities for each token in the generated output, which can be transformed into confidence scores.

Cost Model. Each task $T_i = (m_i, o_i, f_i, \tau_i)$ sends a prompt to model m_i by concatenating the document fraction x_{f_i} and the operation o_i . We denote the number of tokens or *size* as $|\cdot|$, so $|x_{f_i}|$ is the size of the fractional document, and $|o_i|$ is the size of the operation prompt. We let λ_{in} and λ_{cached} denote the cost per new input token and per cached input token, respectively. For

classification tasks, the output cost is typically insignificant and can be ignored, but can be easily incorporated if needed.

The cost to run task T_i on x_{f_i} depends on how much of the input is already cached from previous calls to the same model (as supported by LLM APIs). The cost is:

$$\text{Cost}(T_i, x) = \begin{cases} |x_{f_i}| \lambda_{\text{in}} + |o_i| \lambda_{\text{in}} & \text{if } x_{f_i} \text{ is new} \\ |x_{f_j}| \lambda_{\text{cached}} + (|x_{f_i}| - |x_{f_j}|) \lambda_{\text{in}} + |o_i| \lambda_{\text{in}} & \text{if } x_{f_j} \subseteq x_{f_i} \end{cases}$$

where x_{f_j} is the largest previously processed document prefix that is a subset of x_{f_i} . Throughout, we place the document before the operation in the prompt, maximizing cache utilization of document tokens across multiple tasks using the same model. The total inference cost for document x is then:

$$\text{Cost}(\pi, x) = \sum_{i=1}^{j^*} \text{Cost}(T_i, x)$$

where j^* is the index of the first task to return a confident prediction (i.e., where $\text{Cascade}(\pi, x)$ is determined), or $j^* = k + 1$ if the oracle is called on the original document.

Problem Statement. Given a document collection D , an original operation o_{orig} , a set of classes C , an oracle model m_{oracle} , available proxy models M , and an accuracy target α , we seek to construct a minimal-cost task cascade $\pi = (T_1, T_2, \dots, T_k)$ as follows:

$$\begin{aligned} \min_{\pi} \quad & \sum_{x \in D} \text{Cost}(\pi, x) \\ \text{s.t.} \quad & \Pr \left[\frac{1}{|D|} \sum_{x \in D} \mathbb{I}[\text{Cascade}(\pi, x) = T_{k+1}(x)] \geq \alpha \right] \geq 1 - \delta \end{aligned}$$

where $\mathbb{I}[\cdot]$ is the indicator function and T_{k+1} is the oracle task; i.e., $T_{k+1} = (m_{\text{oracle}}, o_{\text{orig}}, 1, \emptyset)$; so $T_{k+1}(x)$ denotes the oracle's prediction for x . $\delta \in (0, 1)$ is a user-specified upper bound on the probability that the cascade's accuracy on D falls below α .

While we focus on classification with fixed output classes, the task cascade framework can be extended to open-ended map or generation operations, provided there is a reliable automatic evaluator for correctness of an output (i.e., an LLM-as-a-judge [69] and a means to compute confidence scores, e.g., the geometric mean of per-token probabilities as in [23].

2.2 Overview of Our Approach

Our approach constructs a cost-effective task cascade π in an offline phase, using a small, representative sample of the document collection D , which we'll refer to as D_{dev} , the *development set*. The complete procedure is summarized in Algorithm 1. The construction involves: (i) restructuring documents to support efficient fractional processing, and (ii) an agentic loop that iteratively refines a set of candidate tasks and assembles an optimized cascade. We focus on a two-model setting—a cheap proxy model (e.g., GPT-4o-mini) and a more accurate, expensive oracle (e.g., GPT-4o)—as in model cascades, but our approach can generalize to any number of models. An overview of our approach is depicted in Figure 2.

2.2.1 Document Restructuring. We pre-process each document to ensure that the most relevant content for the original operation o_{orig} appears at the beginning. This reordered layout lets us reuse computation: if a longer document fraction is needed later in the cascade, the tokens from shorter prefixes have already been processed and cached by the model, so we only pay to process the

Algorithm 1: Task Cascades Overview

Input: Documents D , original operation o_{orig} , oracle and proxy models
Output: Task cascade π (optionally with statistical guarantees)

```

// Prepare development set
2 Sample  $D_{\text{dev}}$  from  $D$ ;
  // Document Restructuring (Section 4)
3 foreach document  $x \in D_{\text{dev}}$  do
4   Train lightweight classifier to score chunk relevance;
5   Reorder  $x$  to front-load relevant content;
6 end

  // Initialize candidate tasks with original operation at multiple document fractions  $\mathcal{F}$ 
7  $\mathcal{T} \leftarrow \{(m, o_{\text{orig}}, f) : m \in \{\text{proxy, oracle}\}, f \in \mathcal{F}\}$ ;
  // Agentic Loop for Surrogate Discovery
8 for  $n_a$  iterations do
9   // Assemble best cascade from current candidates (Section 3)
   $\pi \leftarrow \text{GreedyCascadeAssembly}(\mathcal{T}, D_{\text{dev}})$ ;
  // Generate new surrogate operations using agent (Section 5)
10   $O_{\text{new}} \leftarrow \text{LLMAgent}(o_{\text{orig}}, \pi)$ ;
  // Add new candidates to task set
   $\mathcal{T} \leftarrow \mathcal{T} \cup \{(m, o, f) : o \in O_{\text{new}}, m \in \{\text{proxy, oracle}\}, f \in \mathcal{F}\}$ ;
11  if no cost improvement then
12    | break;
13  end
14 end
15 return  $\pi$ ;

```

additional new content. We use the oracle model to label regions of each document by relevance to o_{orig} , then train a lightweight classifier to automate reordering for new documents not in D_{dev} .

2.2.2 Agentic Loop for Cascade Construction. To build π , we iteratively grow a set of candidate task configurations $\mathcal{T}$ and repeatedly assemble the lowest-cost cascade that meets the target accuracy α on D_{dev} . Each candidate task configuration is defined by a model, an operation, and a document fraction—that is, the fraction of the reordered document (starting from the beginning) made available to the model. We use a fixed set $\mathcal{F}$ of fractions (e.g., $\mathcal{F} = \{0.1, 0.25, 0.5, 1.0\}$). Initially, we populate the candidate set $\mathcal{T}$ with all combinations of the original operation o_{orig} , both models, and these document fractions. (Confidence thresholds for each task are determined during cascade assembly.)

Each iteration of the agentic loop proceeds as follows. First, we assemble the lowest-cost cascade, π_{best} , that achieves the accuracy target, given the current $\mathcal{T}$. Next, we identify π_{best} ’s “failure cases,” i.e., documents that reach the oracle or are not confidently classified by any task in the cascade. We then prompt an LLM agent to propose new surrogate operations, providing context about these failure cases to inform the agent’s proposals. For each new surrogate, we pair it with all models and document fractions and add the resulting tasks to $\mathcal{T}$. We repeat this loop for a fixed number of rounds or until no further cost improvements are found.

Our approach has three main technical components, and we describe them in the remainder of the paper. First, we present cascade assembly (Section 3), which operates independently of how tasks in the cascade are produced. We then follow the approach to creating candidate tasks as depicted in Figure 2: document restructuring (Section 4) followed by surrogate generation (Section 5). These components are highly modular and can each be customized or replaced.

Table 1. Table of Notation

Symbol	Description
$x \in D$	A document in a collection D
D_{dev}	Development set of documents (to find a cascade)
o_{orig}	Original operation, in natural language
$c \in C$	A class from a predefined set of classes C
α	Accuracy target, where $\alpha \in (0, 1]$
$m; m_{\text{oracle}} \in \mathcal{M}$	Any model; oracle model in the set of models
$p_m(c \mid x_f, o)$	Score produced by m for class c , given doc fraction x_f and operation o
τ^c	Class-specific confidence threshold for class c
(m_i, o_i, f_i)	A task config., defined by m_i , operation o_i , doc fraction f_i
$T_i = (m_i, o_i, f_i, \tau_i)$	A task, comprising a task config. and confidence thresholds τ_i
$\pi = (T_1, T_2, \dots, T_k)$	A task cascade, i.e., ordered sequence of tasks
$\text{Cost}(T_i, x)$	Cost to run task T_i on document x
$ \cdot $	Number of tokens or size
$\lambda_{\text{in}}, \lambda_{\text{cached}}$	Cost per new (cached) input token
$f \in \mathcal{F}$	Set of document fractions to consider for each task
$n_s; n_a$	Number of new surrogate operations per iteration; number of iterations
g	Minimum fraction of D_{dev} classified by a task for inclusion in cascade

3 Cascade Construction

We describe how to construct a cost-efficient cascade from a set of candidate task configurations $\mathcal{T} = \{(m_i, o_i, f_i)\}$, comprising model, operation, and document fraction tuples. Our objective is to assemble an ordered sequence of tasks, with associated confidence thresholds, that minimizes inference cost while meeting a target accuracy α , as formalized in Section 2.1. For this section, we assume $\mathcal{T}$ is fixed; Section 4 will describe how we reorder content within each document, creating fractional documents, and Section 5 will explain how we generate new surrogate operations. Throughout, we assume a fixed set of document fractions per model and operation (original or surrogate).

We first show that a substantially simplified version of cascade assembly is NP-HARD (Section 3.1), motivating our greedy approach. We then detail a three-step procedure for constructing the cascade (Section 3.2): (i) filtering out candidate tasks that cannot meet the accuracy requirement; (ii) greedily assembling an ordered cascade that maintains per-task accuracy; and (iii) adjusting the cascade to meet the overall accuracy guarantee.

3.1 NP-Hardness of Optimal Task Cascade

Constructing an optimal task cascade π on D_{dev} , given task configurations $\mathcal{T} = \{(m_i, o_i, f_i)\}$ is NP-HARD. We show hardness holds even if all $m_i = m$, a single proxy model, and all fractions $f_i = 1$. Our reduction is from the NP-HARD MIN-SUM-SET-COVER problem [18] (MSSC), a variant of the more standard SET-COVER problem. In MSSC, we similarly pick a subset of sets that cover a universe of items; however, the output is an ordering of these sets, where, for each item, we pay a cost based on the earliest set that covers it; specifically, if the i th set in the sequence covers item j , then the cost for covering j is i . Our goal is to minimize the total cost of covering all items.

THEOREM 3.1. *Constructing an optimal task cascade is NP-HARD.*

PROOF. (Sketch) Given an instance of MSSC, (U, S) , where U denotes the items, and S denotes the sets, we generate an instance of cascade assembly as follows. For each item $u \in U$, we create a document d_u . For each set $S_i \in S$, we create a candidate task $T_i = (m, o_i, 1)$, where m is a single proxy model, and o_i is an operation that predicts TRUE with confidence 1 on d_u iff $u \in S_i$ and returns a random answer with confidence 0 otherwise. Therefore, every task in our cascade will have thresholds set to 1.

We design the cost model so that accessing cached document tokens costs 0; each document is cached after the first task that processes it. Thus, the total cost for processing document tokens is constant across all cascades and can be ignored. We set each operation o_i to have cost 1, so running

any task T_i on any d_u incurs cost 1. We set the accuracy target $\alpha = 1$ and assign the oracle infinite cost, so every document must exit the cascade via some task. Under this cost model, processing any d_u through a cascade $(T_{i_1}, T_{i_2}, \dots)$ incurs a cost equal to the index of the first T_{i_j} with confidence 1 on d_u . Hence the cascade cost equals the MSSC objective. $\square$

Feige et al. [18] describe a greedy algorithm for MSSC, which, at each stage, picks the set that covers the most (uncovered) items, and has an approximation factor of 4. Feige et al. demonstrate that further improvements to the approximation factor are difficult. We present an analogous greedy algorithm, adapted to handle varying models, operations, costs, document fractions, and accuracies.

3.2 Assembling a Task Cascade

At a high level, assembling a task cascade from a fixed candidate set of task configurations $\mathcal{T}$ consists of the following steps:

- (1) **Threshold selection and filtering.** For each candidate task individually, we find the smallest confidence thresholds such that at least $g\%$ of D_{dev} are classified by the task, with accuracy at least α . We remove any task for which both criteria are not met for any class.
- (2) **Cascade assembly.** We build a cascade by greedily adding the task that most reduces total inference cost at each step, such that each task achieves the target accuracy on the subset of documents it classifies.
- (3) **Threshold adjustment.** We adjust the thresholds in the assembled cascade to ensure that the overall sequence meets the desired accuracy.

We discuss each of these steps in turn.

3.2.1 Threshold Selection and Filtering. Our candidate set $\mathcal{T}$ may contain hundreds of task configurations, many of which have no chance of meeting the target accuracy, e.g., tasks based on poor surrogate operations. We eliminate these by discarding any task configuration that individually cannot confidently classify enough documents at the required accuracy. For the remaining, we determine confidence thresholds up front; these thresholds are used in subsequent cascade assembly.

We apply the procedure in Algorithm 2: for each candidate task, we examine its predictions on D_{dev} and, for each class, find the lowest confidence threshold such that predictions above this threshold achieve the target accuracy α . If no such threshold exists for a class (i.e., the required accuracy cannot be achieved for any threshold), we set it to be ∞ , ensuring any documents that are predicted to be this class by the task do not exit the cascade. If, across all classes, these thresholds allow the task to classify at least a minimum fraction g of D_{dev} , we retain the task and record its thresholds; otherwise, we discard it. In our implementation, we set g to be 10%; g could also be set using statistical bounds (e.g., from Hoeffding's inequality), but for high accuracy targets, we would require many examples per task (e.g., > 250 for $\alpha = 0.95$ and $\delta = 0.25$). We still ensure statistical guarantees on our eventual cascade, as we will discuss in Section 3.2.3.

3.2.2 Cascade Assembly. After filtering, we build the cascade greedily (detailed pseudocode is in Algorithm 4 in the Appendix). We start with the empty cascade π_0 that invokes the oracle on all documents. At each step, we consider inserting each unused candidate task at the end of π , creating a new candidate cascade π' . For each such candidate cascade π' , we execute it on D_{dev} to compute its cost, and to measure the accuracy of every task in π' on the subset of documents it classifies (i.e., those reaching that task, with the output having confidence above its threshold). We only consider π' where every task, including the newly added one, achieves the target accuracy α on its assigned subset. Among those that meet the per-task accuracy requirement, we select the one with

Algorithm 2: Find Thresholds for a Given Task Config.**Input:** Candidate task config. $T = (m, o, f)$, development set D_{dev} , accuracy target α **Output:** Per-class thresholds $\{\tau^c\}$, or discard T

```

1 Initialize  $\tau^c \leftarrow \infty$  for all  $c$ ; total  $\leftarrow 0$ ;
2 foreach class  $c$  do
3    $P_c \leftarrow$  confidence scores assigned to class  $c$  by  $T$  on all  $x \in D_{\text{dev}}$  where  $c$  is the model's predicted class for  $x$ 
   (i.e.,  $\arg \max_{c'} p_m(c' \mid x_f, o) = c$ ); sorted by confidence;
4   foreach unique  $t$  in  $P_c$  (asc) do
5      $S \leftarrow \{x \in D_{\text{dev}} : \arg \max_{c'} p_m(c' \mid x_f, o) = c, p_m(c \mid x_f, o) \geq t\}$ ;
6     if  $|S| > 0$  and accuracy of assigning  $S$  to  $c \geq \alpha$  then
7        $\tau^c \leftarrow t$ ; total  $+= |S|$ ;
8       break
9     end
10  end
11 end
12 return  $\{\tau^c\}$  if total  $\geq g \cdot |D_{\text{dev}}|$ ; else discard  $T$ ;

```

the lowest total cost. If no such candidate cascade reduces the cost relative to the current cascade, we return the current cascade.

By requiring every task in the cascade to independently satisfy the target accuracy on its assigned documents, we may construct a more conservative cascade than strictly necessary. However, since D_{dev} may be small, a cascade that merely fits the overall accuracy constraint on the development set may not generalize at inference time. Enforcing per-task accuracy thus increases the likelihood that the cascade meets the desired target. Note that for simplicity, this greedy assembly algorithm operates on tasks with fixed, pre-determined confidence thresholds (set in Section 3.2.1). The subsequent threshold adjustment step (Section 3.2.3) is purely a post-processing step to provide statistical accuracy guarantees, and does not further optimize cascade cost.

Algorithm 3: Threshold Adjustment**Input:** Cascade thresholds $\mathcal{T}_0$ and validation set D_V **Output:** Cascade thresholds guaranteed to meet the target accuracy

```

1 for iteration  $i$  until maximum iteration  $\text{max\_iter}$  do
2    $\mathcal{T}_i \leftarrow \{\tau_j + \epsilon_j^i; \forall \tau_j \in \mathcal{T}_0\}$ 
3   if not  $\mathcal{E}(\mathcal{T}_i, D_V)$  then
4     if  $i > 1$  then
5       return  $\mathcal{T}_{i-1}$ ;
6     end
7     else
8       return Not Found;
9     end
10  end
11 end
12 return  $\mathcal{T}_{\text{max\_iter}}$ ;

```

3.2.3 Meeting Accuracy Targets with Guarantees. The cascade π assembled so far is only empirically accurate on the development set and may not achieve the desired accuracy α on new data. To obtain statistical guarantees, we randomly partition the development set D_{dev} into two i.i.d. splits of equal size: a “training” split D_T , used to construct the cascade π , and a “validation” split D_V ,

used to adjust thresholds and yield a final cascade π^* such that $\Pr[\text{Acc}(\pi^*) < \alpha] \leq \delta$. If D_T or D_V are too small or highly variable, the procedure may simply fail to find any cascade meeting the target accuracy under the specified bound; accordingly, our experiments use at least 75 documents in each split. The new task cascade $\pi^* = (T_1^*, \dots, T_k^*)$ is composed of tasks $T_i^* = (m_i, o_i, f_i, \tau_i^*)$ where for all tasks $T_i = (m_i, o_i, f_i, \tau_i)$, m_i , o_i and f_i are kept as is, but the thresholds τ_i are adjusted to τ_i^* . For convenience, we will refer to $\mathbf{t} = \{\tau; \tau \in \tau_i^C, \forall C \forall i\}$ as the set of all thresholds in all tasks in π .

At a high level, determining thresholds that guarantee the target accuracy proceeds as follows. We start by incrementing each threshold in $\mathbf{t}$ by a large non-negative offset $\epsilon_j^{(1)}$, producing an initial, highly conservative set of thresholds. We then iterate through (1) *adjusting* this offset, and (2) *estimating* whether the resulting thresholds are sufficient. In the adjustment phase, at each iteration i , we reduce the offsets to create a new candidate threshold set $\mathbf{t}^{(i)} = \tau_j + \epsilon_j^{(i)} : \tau_j \in \mathbf{t}$, where the adjustment strategy $\epsilon_j^{(i)}$ is chosen to decrease with each step; i.e., $\epsilon_j^{(i+k)} \leq \epsilon_j^{(i)}$ for $k \geq 0$. We will discuss our exact adjustment strategy in Section 3.2.4. In the estimation phase, we run the cascade with $\mathbf{t}^{(i)}$ on the validation split D_V and apply an estimation function $\mathcal{E}(\mathbf{t}^{(i)}, D_V)$. This function returns `True` if the cascade achieves accuracy at least α on D_V with failure probability at most δ , and `False` otherwise. This procedure is presented in Algorithm 3, where the two steps of *adjust* and *estimate* are repeated until a possible maximum number of iterations, and we stop at the first instance where the estimation function determines that the thresholds do not meet the target. If no candidate set passes, we revert to the oracle-only cascade (which never happens in our experiments).

THEOREM 3.2. *For any adjustment strategy ϵ and appropriate estimator $\mathcal{E}$, Algorithm 3 returns a cascade π^* with $\Pr(\text{Acc}(\pi^*) < \alpha) \leq \delta$.*

Theorem 3.2 shows that Algorithm 3 yields the desired accuracy guarantee for any valid estimator and monotonic adjustment schedule. A valid estimator $\mathcal{E}$ must provide a statistical test that, given a candidate set of thresholds and a validation set, returns `True` only if it can certify—based on the observed sample—that the cascade achieves accuracy at least α with failure probability at most δ . This is a classic problem of mean estimation from i.i.d. Bernoulli samples, and any sound concentration bound can be used; e.g., Hoeffding’s inequality. We adopt the estimator from Waudby-Smith and Ramdas [64], which uses both the sample mean and variance to produce tighter bounds—particularly when the variance is small—compared to Hoeffding’s inequality, which relies only on the mean. The proof of Theorem 3.2 and the full specification of $\mathcal{E}$ are given in Appendix A.3 in our technical report [52]. We also account for multiple applications of $\mathcal{E}$ within the adjustment loop, ensuring the total probability of failure remains bounded by δ .

3.2.4 Threshold Adjustment Strategy. We now describe our strategy for adjusting thresholds to ensure the final cascade meets the desired accuracy guarantee. A naive approach is to decrement each threshold by a fixed amount at every iteration (e.g., -0.01 per step). However, LLM-generated confidence scores are often poorly calibrated and heavily concentrated near 1, making uniform decrements ineffective [14, 28, 31]. Instead, for each class c of each task $T = (m, o, f, \tau)$, we collect all confidence scores greater than the initial threshold τ^c assigned by the model to predictions of class c on the training split D_T . We sort these confidence scores in ascending order: $\tau^c < p_1 < p_2 < \dots < p_k$, where p_1 is the lowest confidence just above τ^c and p_k is the highest.

We introduce a *shift index* s that starts at $s_{\max}$ and decrements by 1 at each iteration. At each iteration, for every class, we set its threshold to p_s , corresponding to a more conservative cutoff when s is large, and a less conservative one as s decreases. Specifically, we start with $p_{s_{\max}}$ (the most conservative threshold), and at each subsequent iteration decrement s by 1, thereby shifting the threshold closer to the original τ^c . Once $s = 0$, we use the original threshold τ^c . After updating

all thresholds in this way for an iteration, we evaluate the cascade on D_V and apply the estimator $\mathcal{E}$ to check the accuracy guarantee. If the guarantee is met, we continue; if not, we return the thresholds from the previous shift. Our complete algorithm is given in Appendix A.3 in our technical report [52]. In our implementation, we set $s_{\max} = 5$, but values between 3 and 10 are generally effective.

4 Document Restructuring

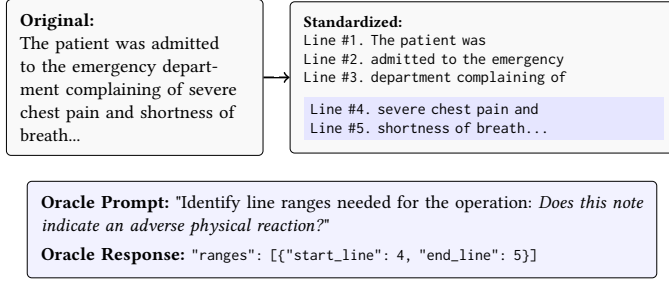


Fig. 3. Standardizing documents. This format allows the oracle to specify relevant content as line ranges.

Tasks in a cascade can process different fractions of each document, allowing models to focus on the most relevant content for the user-defined operation. To reduce cost under our prefix-based cost model (Section 2.1), we reorder each document so that the most relevant text appears first. This restructuring step is independent of the choice of surrogate operations or cascade configuration; any method for ranking document segments by relevance can be used. While document restructuring is optional, skipping it forgoes the prefix-caching discounts that substantially reduce cost.

To restructure documents effectively, we first choose how to segment them. Some user operations require only short spans of text (e.g., a sentence), while others need broader context. We represent each document as a sequence of lines and select a segmentation *granularity*—the number of consecutive lines to treat as a unit. Each document is then divided into contiguous *chunks* of this granularity. Then, we train a lightweight classifier to score and rank chunks in a document by their relevance to the user-defined operation. The following sections describe how we determine the chunk granularity and train the relevance classifier.

Determining Chunk Granularity. To select the chunk granularity for the user-specified operation o_{orig} , we use the oracle LLM and the development set D_{dev} :

- (1) We split each document into lines of 80 characters, each prefixed with a line number (see Figure 3).
- (2) For each document, we ask the oracle LLM to return the minimal set of non-overlapping line ranges needed to answer o_{orig} . The oracle returns these line ranges in a structured format (e.g., "start_line": 5, "end_line": 12).
- (3) We merge overlapping or adjacent line ranges and create a reduced document containing only the union of these lines. We run the oracle LLM on this reduced version and check whether its answer matches the answer on the full document.
- (4) If, across D_{dev} , this reduced document yields the correct answer for at least an α fraction of documents, we move to (5). Otherwise, we expand each range by one line before the start and one line after the end (while staying within document boundaries), merge again, and repeat the check. This process continues until the target accuracy is achieved or a maximum of $e = 3$ expansions.

- (5) We set the chunk granularity to the average length, in lines, of the final merged ranges across all documents in D_{dev} .

Consider the following example for a document x to illustrate the expansion and merging procedure. Suppose the oracle initially selects the ranges $[23, 25]$ and $[28, 30]$, with average length $(3+3)/2 = 3$. If the answers produced by running o_{orig} on the selected lines do not match the answers on the full document for at least an α fraction of documents in D_{dev} , we expand each range for x by one line at both ends, resulting in $[22, 26]$ and $[27, 31]$, with average length $(5+5)/2 = 5$. If another expansion is needed, we obtain $[21, 27]$ and $[26, 32]$, which now overlap and are merged into a single range, $[21, 32]$, with average length 12.

Scoring Chunks by Relevance. Once we determine the segmentation granularity s , we divide each document into contiguous sets of s lines—referred to as *chunks*. Our goal is to train a lightweight *relevance classifier* that predicts whether a chunk contains information required to perform o_{orig} . We first partition D_{dev} into “training” (D_{train}) and “held-out test” (D_{test}) splits for constructing and evaluating the classifier. From the previous step, we already have starting lines of relevant chunks for all documents. We construct an oracle-labeled dataset for training the relevance classifier as follows:

- (1) For each starting line identified by the oracle, we extract the s -line chunk beginning at that line and label it as relevant ($r = 1$). To construct irrelevant examples, we slide a non-overlapping window of s lines from the start of the document, and include a chunk as irrelevant ($r = 0$) only if it does not overlap with any relevant chunk. Chunks that overlap with any relevant chunk are excluded from the irrelevant set.
- (2) For each document in D_{train} and D_{test} , we collect all labeled relevant and irrelevant chunks as described above. The training and test datasets are formed by taking the union of all labeled chunks from each document in D_{train} and D_{test} , respectively.

Each chunk c in D_{train} and D_{test} is embedded using OpenAI’s text-embedding-3-small, producing pairs (embedding(c), r). Because relevant chunks are rare, we upsample the relevant class during training. We fit a logistic regression model with weights initialized to the embedding of o_{orig} to predict the relevance for each chunk. The model is trained to minimize binary cross-entropy loss on the chunks from D_{train} , using the Adam optimizer and early stopping based on F1 score on the chunks from D_{test} .

At inference time, a new document is partitioned into consecutive, non-overlapping chunks of s lines; each chunk is embedded and assigned a probability of relevance by the trained model. Chunks are then sorted by predicted probability (from highest to lowest) and concatenated to yield the reordered document.

5 Surrogate Operation Generation

So far, we have described how to assemble a cost-optimal cascade from a fixed set of candidate task configurations. We now address the question of *how to create this candidate set*. We assume that all documents have already been reordered so that content most relevant to the original operation appears at the beginning (see Section 4 for details), and we only discuss how to generate surrogate operations. At a high level, our approach proceeds as follows. We initialize $\mathcal{T}$ to include all combinations of o_{orig} , each model $m \in \mathcal{M}$, and each fraction $f \in \mathcal{F}$. We then enter a loop that, at each round, (1) assembles the best cascade from the current $\mathcal{T}$, (2) analyzes the failure cases of this cascade, and (3) expands $\mathcal{T}$ using an LLM agent. We next explain what constitutes a valid surrogate operation and how surrogate operations are elicited from the agent, and finally, how the agentic loop is implemented.

Valid Surrogate Operations. In order to generate surrogate operations, we require some notion of validity. A valid surrogate operation is any operation whose outputs form a subset of the original operation’s classes. For classification tasks, a valid surrogate operation is simply any operation whose output space is a subset of the original classes (optionally including a special “none of the above” or -1 label). Suppose the original task is to classify biomedical abstracts into four classes: 0 (Research Article), 1 (Review Article), 2 (Clinical Trial), and 3 (Case Report). A surrogate operation could be: *“If the abstract contains phrases like “randomized trial” or “double-blind,” output 2 (Clinical Trial); otherwise, output -1 .”* Such heuristics enable proxy models to confidently resolve certain examples with simple patterns, even if they cannot perform full multi-class classification.

Eliciting Surrogate Operations from the Agent. Using the aforementioned notion of surrogate operations, whenever we query the LLM agent for surrogates, we include the following context:

- (1) The user-defined operation, o_{orig} .
- (2) “Failure cases,” or examples of documents that are *not* classified by any existing proxy task in the current best cascade (i.e., documents that reach the oracle). To ensure we can fit multiple examples of documents within the agent’s prompt window, we use the oracle model to identify and concatenate only the most relevant snippets from each document—those that support or explain the oracle’s prediction—rather than the full document.
- (3) For each task in the current cascade: summary statistics including whether it was selected into the cascade, its document coverage (i.e., number of documents it resolves), and up to 10 “hard examples,” or borderline cases where the task predicted incorrectly with high confidence, typically those with confidence just above the acceptance threshold.
- (4) Explicit instructions to propose new surrogate operations that (i) are simpler than the original, (ii) output any subset of the original classes, or (iii) target weaknesses in the current cascade. The prompt includes concrete examples and specifies an output format (SURROGATE PROMPT: . . . , RATIONALE: . . .), so we can parse the surrogate operations. See Appendix B in our technical report [52] for details.

Agentic Loop. Now that we have a mechanism for eliciting surrogate operations from the agent, we describe the full agentic loop. At each round: (1) We assemble the best cascade using the current candidate set, via the procedure in Section 3.2; (2) We update the agent’s prompt with the new round’s failure cases and summary statistics; and (3) The agent proposes n_s new surrogate operations. Each surrogate is paired with every model and document fraction, resulting in $|\mathcal{M}| \times |\mathcal{F}| \times n_s$ new candidate tasks, which are added to $\mathcal{T}$. We repeat this process until n_a rounds have completed or no further cost improvement is observed. Parameter choices (e.g., n_s , n_a , document fractions, agent model) are flexible (see Section 2); we $n_s = 3$, $n_a = 3$, and OpenAI’s o1-mini as the agent. See Appendix B in our technical report [52] for the full algorithm.

6 Cost Model for Optimization

In this section, we derive the optimization cost for constructing task cascades and show how hyperparameter choices can control the optimization budget. The total optimization cost comprises three components: document restructuring, surrogate operation generation and evaluation, and agent loop costs. For simplicity, we omit candidate task output token costs (which are constant and negligible—e.g., 1 token for binary classification, at most a few tokens for multi-class tasks). Note that in practice, the actual optimization cost may be lower than our derivation suggests: LLM API providers employ prompt caching, and evaluating multiple candidate tasks on the same documents can lead to cache hits that reduce costs. Our optimization implementation does not explicitly optimize for cache reuse, so the formulas below represent a conservative estimate on the true optimization cost.

Initial labeling and document restructuring. We first run the oracle model once on the full development set to obtain “ground truth” labels for computing surrogate task accuracy: $C_{\text{labels}} = N(L + P)\lambda_o$, where N is the number of development documents, L is the average number of tokens per document, P is the number of prompt tokens per call, and λ_o is the oracle model price per token.

Next, we perform document restructuring to prepare documents for fractional processing. We run the oracle model once more to identify relevant line ranges or chunks of the document. We also generate embeddings for all chunks:

$$C_{\text{doc}} = C_{\text{labels}} + N(L + P)\lambda_o + NL\lambda_{\text{emb}} = N(L + P)(2\lambda_o) + NL\lambda_{\text{emb}}$$

where λ_{emb} is the embedding price per token.

Surrogate operation generation and evaluation.. Over n_a iterations, the agent proposes n_s new surrogate operations per iteration, yielding $n_s n_a$ total surrogates. Each surrogate is paired with each document fraction $f \in \mathcal{F}$ and evaluated on the development set using both the proxy and oracle models. Since we evaluate at $|\mathcal{F}|$ different fractions, the total fractional content processed per document is $S_f = \sum_{i=1}^{|\mathcal{F}|} f_i$ (e.g., for $\mathcal{F} = \{0.15, 0.25, 0.5, 1.0\}$, we have $S_f = 1.90$). With λ_p the proxy model price per token, the cost for all surrogate evaluations is:

$$C_{\text{eval}} = N n_s n_a [LS_f(\lambda_o + \lambda_p) + P|\mathcal{F}|(\lambda_o + \lambda_p)]$$

Agent loop. The LLM agent (o1-mini) generates new surrogates over n_a iterations, consuming T_{agent} input tokens and producing O_{agent} output tokens per round, with prices $\lambda_{o1,\text{in}}$ and $\lambda_{o1,\text{out}}$. Notably, this cost is independent of the development set size N and thus constant: $C_{\text{agent}} = n_a(T_{\text{agent}}\lambda_{o1,\text{in}} + O_{\text{agent}}\lambda_{o1,\text{out}})$.

Total optimization cost. The total optimization cost is $C_{\text{opt}} = C_{\text{doc}} + C_{\text{eval}} + C_{\text{agent}}$. The dominant component is C_{eval} , the cost of evaluating candidate surrogate tasks on the development set. This cost can be reduced by: (1) decreasing n_s or n_a to generate fewer surrogates, (2) using fewer document fractions (smaller $|\mathcal{F}|$), (3) excluding the oracle model when generating candidate tasks (replacing $\lambda_o + \lambda_p$ with λ_p in C_{eval}), or (4) reducing the development set size N , though (4) makes statistical guarantees harder to achieve. In our experiments, we also evaluate a cheaper variant using strategy (3).

7 Evaluation

We design our experiments to answer the following questions:¹

- Q1** Do task cascades reduce inference cost compared to model cascades, and which components of our approach contribute most to these savings?
- Q2** How do cost and accuracy change as target accuracies vary?
- Q3** How consistent are task cascade costs and accuracy across runs?
- Q4** What is the cost of building task cascades?

Across all workloads at a 90% target accuracy, **our approach reduces inference cost by 48.5% on average compared to model cascade baselines (with comparable guarantees), and by 86.2% compared to using the oracle model alone**, highlighting the effectiveness of task cascades (Q1). Every component of our approach (i.e., surrogate operation discovery, document pruning, and cascade ordering) contributes to cost reduction. For Q2, task cascades consistently improve the cost–accuracy tradeoff as the target accuracy varies. The largest gains occur at lower target accuracies, likely because it is easier to identify surrogate operations that satisfy relaxed accuracy targets. On more challenging workloads, the performance gap between model cascades and task cascades narrows at higher accuracy targets. For Q3, we find that task cascades and model cascades

¹Code available at <https://github.com/ucbepic/task-cascades>.

Table 2. Workloads with average word counts and corpus sizes. “# Docs” reports the approximate number of documents in the dataset at release time; actual corpus sizes may have changed in later versions.

Dataset	Description of $\mathcal{D}_{\text{orig}}$	Avg. words	# Docs
AGNEWS	Classify news article summaries into one of four topics: <i>World, Sports, Business, or Science/Tech</i> .	~37	~128k
COURT	Determine if a U.S. Supreme Court opinion reverses the lower-court ruling.	~3.7k	~36k
ENRON	Identify emails sent by C-suite or VP-level executives in the Enron corpus.	~1.5k	~500k
FEVER	Decide whether a natural-language claim is supported by the provided evidence snippets.	~5.1k	185k
GAMES	Determine whether a review praises a different game more than the one being reviewed.	~1.1k	~6.4M
LEGAL	Detect covenants not to sue or IP no-challenge clauses in license agreements.	~8.0k	510
PUBMED	Classify biomedical articles into one of six study types: <i>RCT, Observational, Meta-analysis, Bench/Lab, Computational, or Review</i> .	~3.1k	~133k
WIKI_TALK	Predict whether a Wikipedia Talk-page discussion culminates in an edit revert.	~0.9k	~125k

both exhibit high variance, but task cascades consistently achieve lower mean and median costs. Finally, for **Q4**, we show that optimization costs are negligible at scale, and that even considering optimization cost, task cascades are cheaper than model cascades.

We first describe our experimental setup (Section 7.1). Next, we present results for each research question: cost-effectiveness across variations (Section 7.2.1), performance under varying accuracy targets (Section 7.2.2), consistency across repeated runs (Section 7.2.3), and optimization costs and break-even analysis (Section 7.2.4). Finally, we provide insights from analyzing task cascades, including practical deployment considerations and parameter selection guidance (Section 7.3). We also detail examples of surrogate tasks and the number of tasks in each cascade in Appendix C in our technical report [52].

7.1 Setup

7.1.1 Datasets and Workloads. We evaluate our approach on eight document classification workloads chosen to span a range of domains, document lengths, and task complexities (Table 2). Five are drawn from prior research on LLM-powered data processing systems. To broaden coverage, we introduce three additional tasks from Kaggle that require multi-step or domain-specific reasoning.

From Prior Work. Workloads (i)-(ii) are derived from prior work on model cascades [9, 47], while (iii)-(v) are from work on LLM-powered data processing [40, 51]: (i) AGNEWS [13], from [9], is a four-way classification task with extremely short documents. We include AGNEWS as an example of a simple workload where the proxy already meets the accuracy target, to test whether task cascades can provide additional benefits. (ii) FEVER [58], from [47], is a claim verification task: given a natural-language claim and a set of evidence snippets, the LLM must decide whether the claim is supported. Since each document involves different claims and evidence, reusable filters and surrogate tasks are difficult to learn, making FEVER particularly challenging for our approach. (iii) ENRON [37], from [40], consists of emails from the Enron corpus. The original task in [40] involved fraud detection, but fraud cases are extremely rare in the dataset, so a trivial “always predict no fraud” strategy achieves very high accuracy. To create a more balanced classification problem suitable for evaluating task cascades, we adapt the original fraud detection task to identifying emails sent by C-suite or VP-level executives, which provides a more even class distribution. (iv) LEGAL [25], from [51], tests detection of specific legal clauses in contracts; we convert the original span extraction task into binary classification by focusing on the presence or absence of a single clause type. (v) GAMES [55], from [51], is adapted to a more challenging classification task: determining whether a video game review praises another game more than the one being reviewed.

New Workloads. We introduce three additional workloads from Kaggle and recent NLP datasets. **COURT** [21] involves classifying U.S. Supreme Court opinions and requires multi-step reasoning over lengthy, complex texts. **WIKI_TALK** [6, 12] challenges models to predict whether a Wikipedia Talk-page discussion culminates in an edit revert, capturing dynamics of online discourse. Finally, **PUBMED** [11] is a multi-class classification task on long biomedical articles, testing our approach under both large document size and an expanded label space.

For each workload, we sample 1,000 documents (200 for development, 800 for test), except **LEGAL**, which contains only 509 documents (150 development, 359 test). We sampled datasets to stay within a \$5,000 API budget; running all variants described below across full datasets with multiple accuracy targets and trials would have cost hundreds of thousands of dollars. All workloads represent binary classification tasks except **AGNEWS** and **PUBMED**, which are multi-class. Prompts for all workloads are in Appendix F in our technical report [52].

7.1.2 Baselines. For all experiments in **Q1–Q4**, we compare task cascades against three reference approaches. **Oracle Only** runs the oracle model (GPT-4o) on every document, giving an ideal but costly upper bound. Our second baseline, **2-Model Cascade**, uses the model cascade approach from [47], pairing GPT-4o-mini as proxy with GPT-4o as oracle. This approach is state-of-the-art as it leverages token-level log probabilities from LLM APIs, unlike prior model cascade papers that employ more heavy-weight and less accurate approaches [9, 68]. For each class, we set the proxy threshold to the smallest value on the development set such that the combined accuracy of proxy predictions above the threshold *and* oracle predictions below it meets the target accuracy, aggressively minimizing cost. Third, **2-Model Cascade (+ Guarantees)** augments **2-Model Cascade** with our statistical guarantee procedure (Algorithm 3). Although alternatives such as SUPG [34] could enforce guarantees, we apply the same procedure we do to ensure a controlled comparison.

For **Q2**, which evaluates cost–accuracy tradeoffs across varying accuracy targets, we additionally include **LOTUS (+ Guarantees)** [47]. LOTUS embeds model cascades inside filter operators but provides guarantees on *precision* and *recall*, rather than overall accuracy. Because LOTUS requires separate precision and recall targets instead of a single accuracy target, we cannot directly compare it to our approach at fixed accuracy levels (as in **Q1** and **Q3**). For **Q2**, we implement the LOTUS baseline using their `sem_filter` operator and evaluate it across all combinations of precision and recall targets in $\{0.60, 0.80, 0.90, 0.95\} \times \{0.60, 0.80, 0.90, 0.95\}$ on each binary classification workload.

7.1.3 Variants Compared. Like our baselines, all methods use a single proxy model, along with the oracle model. We evaluate **Task Cascades**, our main approach, as well as **Task Cascades (+ Guarantees)**. We also evaluate **Task Cascades (Lite)**, a lightweight variant that minimizes optimization cost by using the same parameters as **Task Cascades** but evaluating all candidate surrogate tasks only with the proxy model, never the oracle, substantially reducing C_{eval} (as described in Section 6). Note that we distinguish **Task Cascades**, the approach, from task cascades, the outcome of such an approach; throughout this section, **Task Cascades** refers to our approach and is treated as a singular entity (e.g., “**Task Cascades** achieves lower cost”), while a “task cascade” refers to an individual cascade configuration discovered by the approach. We also include different variants of task cascades to isolate the effect of each component of our approach:

- (1) **Surrogate Operation Discovery Variants.** **No Surrogates** disables surrogate discovery, so cascades include only the original user-specified operation at different document fractions. **Single-Iteration** generates all surrogate operations in a single batch, without iterative refinement.
- (2) **Document Pruning Variants.** **No Filtering** disables learned document pruning, so all tasks must process full documents. **Naive RAG Filter** replaces learned pruning with a simple retrieval filter that ranks chunks by cosine similarity to the embedding of the user-defined operation.

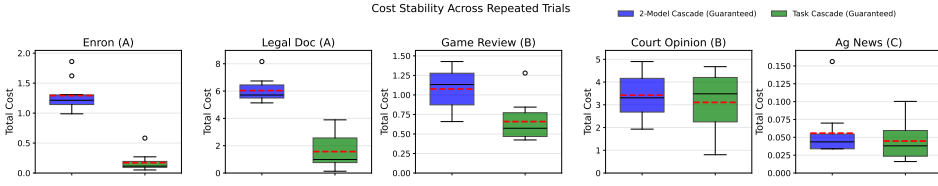


Fig. 4. Cost stability of **2-Model Cascade (+ Guarantees)** and **Task Cascades (+ Guarantees)** across 10 independent runs for five representative workloads. Each box shows the distribution of total inference cost; red dashed lines indicate the median. Both methods exhibit high variance.

(3) **Alternative Cascade Designs.** Rather than using the default greedy strategy, **Selectivity Ordering** constructs the cascade by prioritizing operations with the highest (selectivity -1)/cost ratio, as in prior work on predicate ordering [24], where selectivity is the number of documents not classified by the task and thus passed down the cascade. **Restructure (Top-25%)** applies our learned document restructuring to reorder each document, then allows the proxy to process only the top 25% of the restructured content, while the oracle operates on the full document as usual. This approach isolates the effect of learned restructuring and pruning with a 2-task cascade, without introducing surrogate discovery or having more than 2 tasks. Finally, **RAG + NoSur** applies only naive retrieval-based pruning with no surrogate discovery. This variant allows us to quantify the benefit of simple retrieval and isolate the added value of task cascades.

7.1.4 Metrics and Implementation Details. We report average inference cost in USD (\$) and the fraction of workloads where each method meets the accuracy target α . Most experiments use a fixed $\alpha = 0.9$; in one experiment, we vary α from 0.75 to 0.95. For methods with guarantees, we set failure probability $\delta = 0.25$.

We use OpenAI models: GPT-4o as the oracle model, GPT-4o-mini as the proxy, o1-mini for surrogate generation, and text-embedding-3-small for embeddings. We use PyTorch for the relevance classifier in Section 4. All LLM calls use temperature 0. We consider four document fractions $\mathcal{F} = \{0.1, 0.25, 0.5, 1.0\}$ and run surrogate discovery for three iterations ($n_a = 3$) with three surrogate operations per iteration ($n_s = 5$). These parameter values were held constant across all workloads. While more iterations or surrogates may yield cheaper cascades at higher offline cost, we find these defaults work well across diverse tasks, and run more experiments in Appendix D in our technical report [52] to demonstrate that task cascades are robust to parameter choices.

Inference costs are computed using OpenAI API pricing: \$2.50/1M input tokens and \$10.00/1M output tokens for GPT-4o, \$0.15/1M input tokens and \$0.60/1M output tokens for GPT-4o-mini, with a 50% discount for prefix-cached completions. Embedding costs using (\$0.02/1M tokens) are also included in all reported inference costs; these represent about $0.13\times$ the cost of GPT-4o-mini inference and are negligible compared to overall LLM inference costs. All methods were implemented in Python and run on a 2024 MacBook Air (M4 chip) with OpenAI API calls, costing over \$3,000 USD.

7.2 Results

We now present experimental results addressing the questions outlined above. Table 3 presents the main cost and accuracy results. We reflect on optimization costs in Section 7.3 and leave a detailed discussion of these costs and latencies to Appendix E in our technical report [52].

7.2.1 Q1: Cost Savings and Component Importances. **Task Cascades** reduces inference cost by 41% compared to **2-Model Cascade**, and 48.5% when considering accuracy guarantees. **Task Cascades**

Table 3. Inference results at a 90% accuracy target on a sample of 1k documents per workload (for datasets with more than 1k documents). Each cell shows average accuracy and cost. For all main methods (oracle, 2-model cascades, and task cascades), results are averaged over three trials; variants are single-trial. Baseline costs are shown in USD. Other costs are reported as a multiple of the corresponding 2-Model Cascade variant (plain or +G). The “Avg. Cost” column averages over workloads where the method met the 90% target. Bolded entries mark the lowest cost among methods on each workload.

Method	AGNEWS		COURT		ENRON		FEVER		GAMES		LEGAL		PUBMED		WIKI_TALK		Avg. Cost
Oracle Only	\$0.45		\$11.05		\$6.34		\$13.01		\$3.13		\$10.20		\$8.36		\$2.94		\$6.94
2-Model Cascade	94.1%	\$0.03	89.4%	\$2.81	89.1%	\$0.61	94.4%	\$0.80	89.8%	\$0.78	91.5%	\$3.23	92.5%	\$0.56	91.4%	\$0.18	\$1.13
2-Model Cascades (+G)	94.7%	\$0.04	93.0%	\$4.06	96.2%	\$1.56	96.0%	\$2.22	95.2%	\$1.29	94.8%	\$4.12	93.6%	\$0.83	95.6%	\$0.49	\$1.83
Task Cascades	95.3%	0.66×	88.0%	0.72×	95.5%	0.11×	90.6%	0.54×	92.0%	1.09×	91.3%	0.27×	91.6%	0.84×	91.5%	0.64×	0.59×
Task Cascades (+G)	95.4%	0.71×	90.8%	0.53×	95.6%	0.09×	92.4%	0.44×	89.2%	0.49×	91.2%	0.26×	93.5%	1.23×	92.3%	0.37×	0.52×
Task Cascades (Lite)	94.1%	0.78×	88.6%	0.76×	95.5%	0.14×	90.6%	0.63×	90.2%	0.70×	90.2%	0.41×	91.4%	0.86×	89.9%	0.68×	0.62×
No Surrogates	95.1%	1.43×	94.5%	1.25×	98.0%	0.49×	92.5%	0.86×	90.5%	1.16×	90.0%	0.95×	91.9%	0.79×	96.0%	2.77×	1.21×
Single-Iteration	92.4%	0.56×	93.8%	1.21×	95.9%	0.15×	90.0%	0.30×	89.2%	1.21×	92.8%	1.04×	93.4%	0.76×	90.6%	0.61×	0.66×
No Filtering	93.9%	0.74×	90.0%	1.10×	98.5%	0.96×	95.5%	1.82×	98.4%	1.84×	96.1%	0.93×	95.1%	1.47×	98.4%	3.51×	1.55×
Naive RAG Filter	95.1%	0.95×	90.9%	0.94×	97.2%	0.30×	90.9%	0.42×	85.1%	0.61×	88.9%	0.12×	90.5%	0.88×	89.9%	0.44×	0.65×
Selectivity Ordering	94.2%	1.20×	90.8%	2.33×	98.8%	5.50×	93.1%	5.08×	90.1%	2.32×	92.8%	1.86×	93.8%	8.44×	96.9%	8.81×	4.44×
Restructure (Top-25%)	96.5%	2.67×	93.9%	1.22×	99.4%	0.85×	90.5%	0.37×	97.1%	1.70×	95.3%	1.78×	92.6%	2.38×	98.1%	3.49×	1.81×
RAG + NoSur	96.0%	1.83×	92.1%	1.10×	98.2%	1.73×	91.2%	0.54×	87.9%	1.09×	90.0%	0.77×	90.5%	0.94×	93.4%	1.21×	1.16×

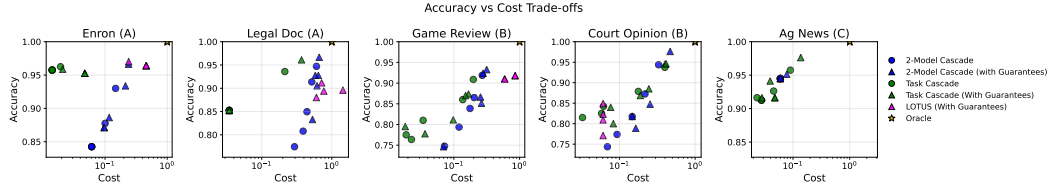


Fig. 5. Accuracy vs. cost trade-offs as the target accuracy varies 75% to 95% (in increments of 5%). Task cascades (green) dominate the Pareto frontier on easy workloads and provide robust gains or new operating points on harder tasks. On simple workloads like AGNEWS, cost improvements appear mainly at lower targets, where the baselines cannot match. LOTUS (+ Guarantees) (pink) only supports binary classification workloads, so we exclude AGNEWS.

(Lite), our lightweight variant with reduced optimization cost, achieves nearly identical savings (0.62× vs. 0.59× on average) while requiring 75% less optimization investment. Savings are most substantial on workloads with significant irrelevant content (e.g., ENRON and LEGAL); inference costs drop by up to 6× relative to 2-Model Cascade.

Meeting the Accuracy Target. Task Cascades and Task Cascades (+ Guarantees) each miss the 90% accuracy target on one workload by < 1%, while 2-Model Cascade fails on 3 of 8 workloads. When accuracy guarantees are required (+ Guarantees variants), both approaches construct cascades using only half the available data, reserving the remainder for threshold adjustment (Algorithm 3), which increases costs by 1.43× on average for Task Cascades.

Component Contributions. All three core components—greedy task ordering, surrogate operation discovery, document pruning—are essential for robust performance. Selectivity Ordering performs worst (7.5× worse than Task Cascades). Single-Iteration (1.13× worse than Task Cascades) demonstrates that iterative refinement contributes to cost reduction beyond single-round surrogate generation. Surprisingly, both No Surrogates (1.21×) and No Filtering (1.55×) are more expensive than the 2-Model baseline. Without surrogates, the proxy is limited to the user-defined operation and provides savings on only half the workloads; on remaining tasks, the entire document must be processed, negating cost benefits. Without document pruning, each surrogate must operate on the entire document, making each surrogate operation as expensive as the first stage of 2-Model Cascade. Similarly, Restructure (Top-25%) fails to improve cost on 6 of 8 workloads: when the proxy

receives too little context to perform the original operation, it cannot confidently resolve documents, escalating more to the expensive oracle. Similarly, **Restructure (Top-25%)** fails to improve cost on 6 of 8 workloads: when the proxy receives too little context to perform the original operation, it cannot confidently resolve documents, escalating more to the expensive oracle. Across workloads, simpler variants can fail in significant ways: **No Filtering** is up to $8\times$ worse than **Task Cascades**, **No Surrogates** up to $4\times$ worse, and the maximum degradation across all variants averages $12\times$ worse than **Task Cascades**. By combining multiple techniques, **Task Cascades** achieves robust cost savings across diverse workloads.

Comparison with Retrieval-Based Pruning. Our learned document pruning outperforms simple retrieval alternatives. **Naive RAG Filter** achieves the second-lowest cost on average, but **Task Cascades** reduces cost by an additional 11%. More notably, **RAG + NoSur** (retrieval-based pruning without surrogate discovery) costs $1.16\times$ more than **2-Model Cascade** on average: naive document reordering can make restructured documents harder for the proxy to process, escalating more documents to the expensive oracle. Adding surrogate discovery (**Naive RAG Filter**) addresses this degradation, making it cheaper than **2-Model Cascade** on all workloads.

Accuracy-Cost Tradeoffs at Fixed Target. At the fixed 90% accuracy target, task cascades explore a broader configuration space, enabling more precise cost-accuracy tradeoffs. In some cases, task cascades identify plans that are both cheaper *and* more accurate than baselines: on ENRON, **Task Cascades** achieves 95.5% accuracy at $0.11\times$ cost vs. **2-Model Cascade**'s 89.1% accuracy at $1.0\times$ cost. In other cases, task cascades find cheaper plans closer to the target: on FEVER, **Task Cascades** achieves 90.6% accuracy at $0.54\times$ cost vs. **2-Model Cascade**'s 94.4% accuracy at $1.0\times$ cost. Overall, **Task Cascades** can search over more configurations and, many times, land closer to the target accuracy.

7.2.2 Q2: Accuracy–Cost Tradeoffs Across Targets. To evaluate performance across different accuracy requirements, we vary the target accuracy from 75% to 95% in 5% increments, using the same workload groups as in Section 7.2.3. Figure 5 shows the resulting accuracy–cost tradeoff. We include **LOTUS (+ Guarantees)** on binary classification workloads (all except AGNEWS), using the setup described in Section 7.1. Overall, **Task Cascades** reduces inference costs and extends the Pareto frontier across both easy and hard tasks, compared to all baselines. **LOTUS (+ Guarantees)** performs comparably to **2-Model Cascade (+ Guarantees)** on all workloads. On easier Group A tasks (ENRON, LEGAL), **Task Cascades** consistently dominates the Pareto frontier, matching or exceeding model cascade accuracy at much lower cost for any target. For more challenging Group B workloads (GAMES, COURT), **Task Cascades** provides the greatest cost savings at lower target accuracies (75%–85%), because the agent can choose from a larger set of surrogate operations that meet these relaxed accuracy targets. As the target accuracy increases, fewer surrogates will satisfy the stricter requirements, so cost advantages diminish—but **Task Cascades** still often matches or outperforms the baseline, perhaps due to document pruning. On Group C (AGNEWS), where the proxy already achieves high accuracy, there is little room for improvement: **Task Cascades** performs similarly to the baseline at high targets but can identify lower-cost solutions at lower accuracy targets.

7.2.3 Q3: Consistency and Variance Across Repeated Trials. We next evaluate the consistency of task cascade outcomes across repeated runs, given the inherent stochasticity of LLM-based surrogate discovery. To provide a representative analysis, while managing compute and cost constraints, we select five workloads spanning a range of document lengths and task complexities. We organize these into three groups based on observed outcomes in Table 3: Group A (ENRON and LEGAL), where task cascades deliver the largest cost reductions; Group B (GAMES and COURT), where it is most

difficult to find cheap and accurate cascades; and Group C (AGNEWS), where the proxy model alone is already highly accurate and documents are short, impeding savings from document pruning techniques.

We next evaluate the consistency of task cascade outcomes across repeated runs, given the inherent stochasticity of LLM-based surrogate discovery. To provide a representative analysis while managing compute and cost constraints, we select five workloads spanning a range of document lengths and task complexities. We group them by observed difficulty in Table 3: Group A (ENRON, LEGAL) where task cascades achieve the largest cost reductions; Group B (GAMES, COURT) where it is hardest to find cheap, accurate cascades; and Group C (AGNEWS) where documents are short and the proxy model is already highly accurate. As shown in Figure 4, we evaluate both **2-Model Cascade (+ Guarantees)** and **Task Cascades (+ Guarantees)** across ten independent runs per workload. On Group A workloads, **Task Cascades (+ Guarantees)** reduces average cost by 86.5% on ENRON and 74.1% on LEGAL, and even the *most expensive* task cascade run remains *cheaper than the cheapest* 2-model cascade. On Group B workloads, both methods show wider variance due to increased task difficulty, but the *75th percentile* task cascade cost is still *lower than the 25th percentile* of the baseline cost. On Group C (AGNEWS), where the baseline is already near-optimal, **Task Cascades (+ Guarantees)** match its cost and variance.

Both methods exhibit notable run-to-run variance. **2-Model Cascade (+ Guarantees)** also shows high variability—likely because the small development set size and data-splitting procedure for threshold adjustment introduce uncertainty in cascade construction. **Task Cascades (+ Guarantees)** consistently achieves lower mean and median cost across all workloads, demonstrating that variability does not undermine its advantage.

7.2.4 Q4: Optimization Costs. So far, we have focused on inference costs, as these dominate in production-scale deployments, where offline optimization costs are offset by online gains. To examine when this investment “pays off”, we report optimization costs, break-even points (when total cost becomes lower than running **Oracle Only**), and estimated costs for processing 1 million documents (selected to showcase the impact at scale). Full details are provided in Appendix E, described in Table 4 in our technical report [52].

At 1M documents, optimization costs become negligible and task cascades provide substantial savings. **Task Cascades (Lite)**—our lightweight variant, described in Section 7.1—provides cost savings over **2-Model Cascade** on all workloads, with average savings of 37.5% and up to 86% on ENRON. **Task Cascades** provides savings on 7 of 8 workloads (average 36%, up to 87% on ENRON); on GAMES, where **Task Cascades** does not provide savings, **2-Model Cascade** fails to meet the 90% accuracy target (Table 3).

These savings are achieved even though **Task Cascades** incurs higher optimization costs than **2-Model Cascade** ($7.9\times$ – $27.0\times$ across workloads, mean $11.4\times$) and takes longer to build ($\sim 14\times$ optimization latency). Our lightweight **Task Cascades (Lite)** variant costs $2.2\times$ – $4.7\times$ the **2-Model Cascade** baseline (mean $2.9\times$) while preserving 95% of inference savings. Optimization cost is recovered after processing 2,986 documents on average for **Task Cascades**, while **Task Cascades (Lite)** requires only 678—modest thresholds relative to our workload sizes (median 129k documents, as in Table 2). Inference latencies are within the same order of magnitude (Table 5 in Appendix E in our technical report [52]).

In summary, while task cascades require higher upfront optimization investment, this cost quickly becomes negligible at realistic scales, demonstrating that task cascades are practical and cost-effective for production deployments.

7.3 Insights from Analyzing Task Cascades

Choosing Variants of the Task Cascade Approach. The **Task Cascades** approach provides the most consistent cost reductions across workloads, making it a reliable default. Simpler variants may be preferable in specific scenarios: **Naive RAG Filter** when engineering effort is limited and documents have clearly identifiable relevant sections; **No Surrogates** when the original operation is not very difficult and the proxy can accurately perform it; and variants without filtering when documents are very short (e.g., AGNEWS). The **Single-Iteration** variant is particularly promising, achieving strong performance with reduced offline cost and likely to improve as LLM agents improve in their ability to reason and handle complex tasks. While individual variants occasionally outperform the full approach, **Task Cascades** is 20% cheaper than the next best variant on average, making it the recommended starting point.

Parameter Selection Guidance. Task cascades depend on three hyperparameters: the number of surrogate operations per iteration (n_s), refinement iterations (n_a), and document fraction sets ($\mathcal{F}$). We conduct sensitivity analysis on two representative workloads, varying $n_s \in \{3, 5, 10\}$, $n_a \in \{1, 2, 3\}$, and testing different $\mathcal{F}$ configurations (detailed results in Appendix D in our technical report [52]). Task cascades consistently outperform baselines across all configurations, demonstrating robustness to parameter choices. Based on these experiments, we recommend: $n_s \geq 3$, $n_a \geq 1$, and document fractions spanning from small (0.1 or 0.25) to full (1.0). We find that n_s has a stronger effect on cost reduction than n_a , so practitioners should prioritize exploring diverse surrogates over extensive refinement.

8 Related Work

We cover three areas: LLM-powered data processing, cost-efficient LLM execution, and query rewriting with LLMs.

LLM-Powered Data Processing. Modern data systems increasingly integrate LLMs as first-class operators for natural language extraction and transformation. Industrial systems like Databricks, DuckDB, Snowflake, and Google’s AlloyDB support LLM-powered *filter* and *map* functions in SQL [16, 20, 54, 65]. LOTUS [47], ThalamusDB [29], and Aryn [1] provide pandas-based, SQL, and Spark-like interfaces respectively. Palimpzest [40], DocETL [51], and AOP [63] offer custom DSLs for LLM-powered data processing. However, these systems typically invoke LLMs on each document independently, resulting in high inference costs. Some systems focus on supporting LLM-powered data processing for particular query or data modalities [3, 39, 43, 59].

Cost Optimization for LLM Inference. Strategies for reducing LLM inference costs include cost-based optimizers, specialized models, caching, ensembling, and model cascades. ABACUS [50] introduces a Cascades-style optimizer [22] over different “physical” implementations of common LLM-powered operators, while ELEET [60] replaces LLM operators with trainable models. Other work reorders LLM calls for KV-cache reuse [42]. Some methods profile multiple LLMs and aggregate their outputs to reduce cost. For example, ThriftLLM [26] and LLMBlender [27], instead of using only the output of the last processed stage as in model cascades, ensemble predictions from all evaluated models. However, in the primary setting we study—where cascade models differ substantially in cost and quality—there is no additional value to be gained by ensembling the oracle’s output (when available) with that of the proxy. SpareLLM [30] profiles a number of LLMs for a given family of tasks, but ultimately selects a single model to handle the remaining tasks, which in our case boils down to a cascade with a single step, either a proxy model or an oracle, for all items, independent of confidence. Our work builds on model cascades [2, 5, 9, 32–34, 44, 47, 62, 68], generalizing them by varying not only the model but also the operation and document fraction at each stage.

Query Rewriting with LLMs. LLM-powered query rewriting has proven effective in retrieval, question-answering [4, 10, 45, 48], and traditional query optimization [15, 38, 56, 70]. DocETL [51] decomposes complex tasks into logically equivalent subtasks for accuracy. In contrast, we rewrite tasks for cheaper, possibly-incorrect operations, dramatically expanding the rewrite space. Task cascades may not resemble logical decompositions—the same operation can appear multiple times at different document fractions.

Our work also builds on approximate query processing (AQP) [7, 19]. While tailored for unstructured data and LLMs, our components mirror classical AQP: surrogate operations resemble approximate predicates [53], document fractions leverage sampling, and we use concentration bounds for accuracy guarantees [8, 34].

9 Conclusion

We introduce *task cascades*, a generalization of the model cascades framework for efficient LLM-powered unstructured text processing. Task cascades vary the model, operation, and document fraction at each stage to minimize inference cost while meeting accuracy targets. Across eight real-world workloads, task cascades reduce inference cost by 48.5% on average compared to model cascades (with comparable guarantees) and by 86.2% relative to oracle-only inference, with ablations confirming the importance of all components of our approach. Overall, task cascades provide a practical and extensible solution for scalable unstructured data analysis.

Acknowledgments

We acknowledge support from grants DGE-2243822, IIS-2129008, IIS-1940759, and IIS-1940757 awarded by the National Science Foundation, funds from the State of California, an NDSEG fellowship, funds from the Alfred P. Sloan Foundation, as well as EPIC lab sponsors: Adobe, Google, G-Research, Microsoft, PromptQL, Sigma Computing, Bridgewater, and Snowflake. Compute credits were provided by Azure, Modal, NSF (via NAIRR), and OpenAI.

References

- [1] Eric Anderson, Jonathan Fritz, Austin Lee, Bohou Li, Mark Lindblad, Henry Lindeman, Alex Meyer, Parth Parmar, Tanvi Ranade, Mehul A. Shah, Benjamin Sowell, Dan Tecuci, Vinayak Thapliyal, and Matt Welsh. 2025. The Design of an LLM-powered Unstructured Analytics System. In *Proceedings of the Conference on Innovative Data Systems Research (CIDR)*. <https://mail.vldb.org/cidrdb/papers/2025/p13-anderson.pdf>
- [2] Michael R Anderson, Michael Cafarella, German Ros, and Thomas F Wenisch. 2019. Physical representation-based predicate optimization for a visual analytics database. In *2019 IEEE 35th International Conference on Data Engineering (ICDE)*. IEEE, 1466–1477.
- [3] Simran Arora, Brandon Yang, Sabri Eyuboglu, Avani Narayan, Andrew Hojel, Immanuel Trummer, and Christopher Ré. 2023. Language Models Enable Simple Systems for Generating Structured Views of Heterogeneous Data Lakes. *Proceedings of the VLDB Endowment* 17, 2 (2023), 92–105.
- [4] Muhammad Imam Luthfi Balaka, David Alexander, Qiming Wang, Yue Gong, Adila Krisnadhi, and Raul Castro Fernandez. 2025. Pneuma: Leveraging llms for tabular data representation and retrieval in an end-to-end system. *Proceedings of the ACM on Management of Data* 3, 3 (2025), 1–28.
- [5] Jiashen Cao, Karan Sarkar, Ramyad Hadidi, Joy Arulraj, and Hyesoon Kim. 2022. Figo: Fine-grained query optimization in video analytics. In *Proceedings of the 2022 International conference on management of data*. 559–572.
- [6] Jonathan P Chang, Caleb Chiam, Liye Fu, Andrew Wang, Justine Zhang, and Cristian Danescu-Niculescu-Mizil. 2020. ConvoKit: A Toolkit for the Analysis of Conversations. In *Proceedings of the 21th Annual Meeting of the Special Interest Group on Discourse and Dialogue*. 57–60.
- [7] Surajit Chaudhuri, Bolin Ding, and Srikanth Kandula. 2017. Approximate query processing: No silver bullet. In *Proceedings of the 2017 ACM International Conference on Management of Data*. 511–519.
- [8] Surajit Chaudhuri, Rajeev Motwani, and Vivek Narasayya. 1998. Random sampling for histogram construction: How much is enough? *ACM SIGMOD Record* 27, 2 (1998), 436–447.
- [9] Lingjiao Chen, Matei Zaharia, and James Zou. 2024. FrugalGPT: How to Use Large Language Models While Reducing Cost and Improving Performance. *Transactions on Machine Learning Research* (2024).

- [10] Peter Baile Chen, Yi Zhang, Mike Cafarella, and Dan Roth. [n.d.]. Can we Retrieve Everything All at Once? ARM: An Alignment-Oriented LLM-based Retrieval Method. In *The 4th Table Representation Learning Workshop at ACL 2025*.
- [11] Arman Cohan, Franck Dernoncourt, Doo Soon Kim, Trung Bui, Seokhwan Kim, Walter Chang, and Nazli Goharian. 2018. A Discourse-Aware Attention Model for Abstractive Summarization of Long Documents. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 2 (Short Papers)*. Association for Computational Linguistics, New Orleans, Louisiana, 615–621. <https://doi.org/10.18653/v1/N18-2097>
- [12] Cristian Danescu-Niculescu-Mizil, Lillian Lee, Bo Pang, and Jon Kleinberg. 2012. Echoes of power: Language effects and power differences in social interaction. In *Proceedings of the 21st international conference on World Wide Web*. 699–708.
- [13] Gianna M Del Corso, Antonio Gulli, and Francesco Romani. 2005. Ranking a stream of news. In *Proceedings of the 14th international conference on World Wide Web*. 97–106.
- [14] Shrey Desai and Greg Durrett. 2020. Calibration of Pre-trained Transformers. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*. 295–302.
- [15] Sriram Dharwada, Himanshu Devrani, Jayant Haritsa, and Harish Doraiswamy. 2025. Query rewriting via llms. *arXiv preprint arXiv:2502.12918* (2025).
- [16] Till Döhmen. 2024. Introducing the prompt() Function: Use the Power of LLMs with SQL! [urlhttps://motherduck.com/blog/sql-llm-prompt-function-gpt-models/](https://motherduck.com/blog/sql-llm-prompt-function-gpt-models/). Accessed: 2025-06-22.
- [17] Yu Fan, Jingwei Ni, Jakob Merane, Etienne Salimbeni, Yang Tian, Yoan Hermstrüwer, Yinya Huang, Mubashara Akhtar, Florian Geering, Oliver Dreyer, et al. 2025. LEXam: Benchmarking Legal Reasoning on 340 Law Exams. *arXiv preprint arXiv:2505.12864* (2025).
- [18] Uriel Feige, László Lovász, and Prasad Tetali. 2004. Approximating min sum set cover. *Algorithmica* 40 (2004), 219–234.
- [19] Minos N Garofalakis and Phillip B Gibbons. 2001. Approximate Query Processing: Taming the TeraBytes.. In *VLDB*, Vol. 10. 645927–672356.
- [20] Google Cloud. 2025. Perform intelligent SQL queries using AlloyDB AI query engine. <https://cloud.google.com/alloydb/docs/ai/evaluate-semantic-queries-ai-operators>. Accessed: 2025-06-22; Last updated: 2025-06-11.
- [21] gqfiddler. 2022. SCOTUS Opinions. <https://www.kaggle.com/datasets/gqfiddler/scotus-opinions>. Accessed: 2025-07-16.
- [22] Goetz Graefe. 1995. The Cascades Framework for Query Optimization. *IEEE Data(base) Engineering Bulletin* 18 (1995), 19–29. <https://api.semanticscholar.org/CorpusID:260706023>
- [23] Neha Gupta, Harikrishna Narasimhan, Wittawat Jitkrittum, Ankit Singh Rawat, Aditya Krishna Menon, and Sanjiv Kumar. [n.d.]. Language Model Cascades: Token-Level Uncertainty And Beyond. In *The Twelfth International Conference on Learning Representations*.
- [24] Joseph M. Hellerstein and Michael Stonebraker. 1993. Predicate migration: optimizing queries with expensive predicates. In *Proceedings of the 1993 ACM SIGMOD International Conference on Management of Data (Washington, D.C., USA) (SIGMOD '93)*. Association for Computing Machinery, New York, NY, USA, 267–276. <https://doi.org/10.1145/170035.170078>
- [25] Dan Hendrycks, Collin Burns, Anya Chen, and Spencer Ball. 2021. CUAD: An Expert-Annotated NLP Dataset for Legal Contract Review. *NeurIPS* (2021).
- [26] Keke Huang, Yimin Shi, Dujian Ding, Yifei Li, Yang Fei, Laks Lakshmanan, and Xiaokui Xiao. 2025. ThriftLLM: On Cost-Effective Selection of Large Language Models for Classification Queries. *arXiv preprint arXiv:2501.04901* (2025).
- [27] Dongfu Jiang, Xiang Ren, and Bill Yuchen Lin. 2023. LLM-Blender: Ensembling Large Language Models with Pairwise Ranking and Generative Fusion. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 14165–14178.
- [28] Zhengbao Jiang, Jun Araki, Haibo Ding, and Graham Neubig. 2021. How can we know when language models know? on the calibration of language models for question answering. *Transactions of the Association for Computational Linguistics* 9 (2021), 962–977.
- [29] Saehan Jo and Immanuel Trummer. 2024. Thalamusdb: Approximate query processing on multi-modal data. *Proceedings of the ACM on Management of Data* 2, 3 (2024), 1–26.
- [30] Saehan Jo and Immanuel Trummer. 2025. SpareLLM: Automatically Selecting Task-Specific Minimum-Cost Large Language Models under Equivalence Constraint. *Proc. ACM Manag. Data* 3, 3, Article 219 (June 2025), 26 pages. <https://doi.org/10.1145/3725356>
- [31] Saurav Kadavath, Tom Conerly, Amanda Askell, Tom Henighan, Dawn Drain, Ethan Perez, Nicholas Schiefer, Zac Hatfield-Dodds, Nova DasSarma, Eli Tran-Johnson, et al. 2022. Language models (mostly) know what they know. *arXiv preprint arXiv:2207.05221* (2022).
- [32] Daniel Kang, Peter Bailis, and Matei Zaharia. [n.d.]. Blazelt: Optimizing Declarative Aggregation and Limit Queries for Neural Network-Based Video Analytics. *Proceedings of the VLDB Endowment* 13, 4 ([n.d.]).

- [33] Daniel Kang, John Emmons, Firas Abuzaid, Peter Bailis, and Matei Zaharia. 2017. NoScope: Optimizing Neural Network Queries over Video at Scale. *Proceedings of the VLDB Endowment* 10, 11 (2017).
- [34] Daniel Kang, Edward Gan, Peter Bailis, Tatsunori Hashimoto, and Matei Zaharia. 2020. Approximate selection with guarantees using proxies. *Proceedings of the VLDB Endowment* 13, 12 (2020), 1990–2003.
- [35] Daniel Kang and John Guibas. 2021. Accelerating Approximate Aggregation Queries with Expensive Predicates. *Proc. VLDB Endow.* 14 (2021), 2341–2354. <https://api.semanticscholar.org/CorpusID:237012342>
- [36] Daniel Kang, John Guibas, Peter D Bailis, Tatsunori Hashimoto, and Matei Zaharia. 2022. Tasti: Semantic indexes for machine learning-based queries over unstructured data. In *Proceedings of the 2022 International Conference on Management of Data*. 1934–1947.
- [37] Bryan Klimt and Yiming Yang. 2004. The enron corpus: A new dataset for email classification research. In *European conference on machine learning*. Springer, 217–226.
- [38] Zhao Donghui Li, Haitao Yuan, Huiming Wang, Gao Cong, and Lidong Bing. 2024. LLM-R2: A Large Language Model Enhanced Rule-Based Rewrite System for Boosting Query Efficiency. *Proceedings of the VLDB Endowment* 18, 1 (2024), 53–65.
- [39] Yiming Lin, Madelon Hulsebos, Ruiying Ma, Shreya Shankar, Sepanta Zeighami, Aditya G Parameswaran, and Eugene Wu. 2025. Querying Templatized Document Collections with Large Language Models. In *2025 IEEE 41st International Conference on Data Engineering (ICDE)*. IEEE Computer Society, 2422–2435.
- [40] Chunwei Liu, Matthew Russo, Michael Cafarella, Lei Cao, Peter Baile Chen, Zui Chen, Michael Franklin, Tim Kraska, Samuel Madden, Rana Shahout, et al. 2025. Palimpsest: Optimizing ai-powered analytics with declarative query processing. In *Proceedings of the Conference on Innovative Database Research (CIDR)*. 2.
- [41] Nelson F Liu, Kevin Lin, John Hewitt, Ashwin Paranjape, Michele Bevilacqua, Fabio Petroni, and Percy Liang. 2024. Lost in the middle: How language models use long contexts. *Transactions of the Association for Computational Linguistics* 12 (2024), 157–173.
- [42] Shu Liu, Asim Biswal, Amog Kamsetty, Audrey Cheng, Luis Gaspar Schroeder, Liana Patel, Shiyi Cao, Xiangxi Mo, Ion Stoica, Joseph E Gonzalez, et al. 2025. Optimizing LLM Queries in Relational Data Analytics Workloads. In *Eighth Conference on Machine Learning and Systems*.
- [43] Duo Lu, Siming Feng, Jonathan Zhou, Franco Solleza, Malte Schwarzkopf, and Uğur Çetintemel. 2025. VectraFlow: Integrating Vectors into Stream Processing. *CIDR*.
- [44] Yao Lu, Aakanksha Chowdhery, Srikanth Kandula, and Surajit Chaudhuri. 2018. Accelerating machine learning inference with probabilistic predicates. In *Proceedings of the 2018 International Conference on Management of Data*. 1493–1508.
- [45] Xinbei Ma, Yeyun Gong, Pengcheng He, Hai Zhao, and Nan Duan. 2023. Query rewriting in retrieval-augmented large language models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. 5303–5315.
- [46] Matteo Magnini, Gianluca Aguzzi, and Sara Montagna. 2025. Open-source small language models for personal medical assistant chatbots. *Intelligence-Based Medicine* 11 (2025), 100197.
- [47] Liana Patel, Siddharth Jha, Melissa Pan, Harshit Gupta, Parth Asawa, Carlos Guestrin, and Matei Zaharia. 2025. Semantic Operators and Their Optimization: Enabling LLM-Based Data Processing with Accuracy Guarantees in LOTUS. *Proceedings of the VLDB Endowment* 18, 11 (2025), 4171–4184.
- [48] Wenjun Peng, Guiyang Li, Yue Jiang, Zilong Wang, Dan Ou, Xiaoyi Zeng, Derong Xu, Tong Xu, and Enhong Chen. 2024. Large language model based long-tail query rewriting in taobao search. In *Companion Proceedings of the ACM Web Conference 2024*. 20–28.
- [49] Matthew Russo, Tatsunori Hashimoto, Daniel Kang, Yi Sun, and Matei Zaharia. 2023. Accelerating Aggregation Queries on Unstructured Streams of Data. *Proceedings of the VLDB Endowment* 16, 11 (2023), 2897–2910.
- [50] Matthew Russo, Sivaprasad Sudhir, Gerardo Vítagliano, Chunwei Liu, Tim Kraska, Samuel Madden, and Michael Cafarella. 2025. Abacus: A Cost-Based Optimizer for Semantic Operator Systems. *arXiv preprint arXiv:2505.14661* (2025).
- [51] Shreya Shankar, Tristan Chambers, Tarak Shah, Aditya G Parameswaran, and Eugene Wu. 2025. DocETL: Agentic Query Rewriting and Evaluation for Complex Document Processing. *Proceedings of the VLDB Endowment* 18, 9 (2025), 3035–3048.
- [52] Shreya Shankar, Sepanta Zeighami, and Aditya Parameswaran. 2026. Task Cascades for Efficient Unstructured Data Processing. *arXiv:2601.05536 [cs.DB]* <https://arxiv.org/abs/2601.05536>
- [53] N. Shivakumar, H. Garcia-Molina, and C. Chekuri. 1998. Filtering with approximate predicates. In *24rd International Conference on Very Large Data Bases (VLDB 1998)*. <http://ilpubs.stanford.edu:8090/351/>
- [54] Snowflake. 2025. Introducing Cortex AISQL: Reimagining SQL into AI Query Language for Multimodal Data. <https://www.snowflake.com/en/blog/ai-sql-query-language/>. Accessed: July 15, 2025.

- [55] Antoni Sobkowicz and Wojciech Stokowiec. 2016. Steam review dataset-new, large scale sentiment dataset. In *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC 2016) Workshop Emotion and Sentiment Analysis*. 55–58.
- [56] Yuyang Song, Hanxu Yan, Jiale Lao, Yibo Wang, Yufei Li, Yuanchun Zhou, Jianguo Wang, and Mingjie Tang. 2025. QUITE: A Query Rewrite System Beyond Rules with LLM Agents. *arXiv preprint arXiv:2506.07675* (2025).
- [57] Sijun Tan, Xiuyu Li, Shishir G Patil, Ziyang Wu, Tianjun Zhang, Kurt Keutzer, Joseph Gonzalez, and Raluca Popa. 2024. LLoCO: Learning Long Contexts Offline. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*. 17605–17621.
- [58] James Thorne, Andreas Vlachos, Christos Christodoulopoulos, and Arpit Mittal. 2018. FEVER: a Large-scale Dataset for Fact Extraction and VERification. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*. 809–819.
- [59] Matthias Urban and Carsten Binnig. 2024. CAESURA: Language Models as Multi-Modal Query Planners. In *14th Conference on Innovative Data Systems Research, CIDR 2024, Chaminade, CA, USA, January 14-17, 2024*. www.cidrdb.org. <https://www.cidrdb.org/cidr2024/papers/p14-urban.pdf>
- [60] Matthias Urban and Carsten Binnig. 2024. Eleet: Efficient learned query execution over text and tables. *Proceedings of the VLDB Endowment* 17, 13 (2024), 4867–4880.
- [61] Vals AI, Inc. 2024. Vals AI. <https://www.vals.ai/>. Accessed: 2025-06-22.
- [62] Paul Viola and Michael Jones. 2001. Rapid object detection using a boosted cascade of simple features. In *Proceedings of the 2001 IEEE computer society conference on computer vision and pattern recognition. CVPR 2001*, Vol. 1. Ieee, I–I.
- [63] Jiayi Wang and Guoliang Li. 2025. Aop: Automated and interactive llm pipeline orchestration for answering complex queries. CIDR.
- [64] Ian Waudby-Smith and Aaditya Ramdas. 2024. Estimating means of bounded random variables by betting. *Journal of the Royal Statistical Society Series B: Statistical Methodology* 86, 1 (2024), 1–27.
- [65] Patrick Wendell, Eric Peter, Nicolas Pelaez, Jianwei Xie, Vinny Vijeyakumaar, Linhong Liu, and Shitao Li. 2023. Introducing AI Functions: Integrating Large Language Models with Databricks SQL. <https://www.databricks.com/blog/2023/04/18/introducing-ai-functions-integrating-large-language-models-databricks-sql.html>. Accessed: 2025-06-22.
- [66] Xinyi Wu, Yifei Wang, Stefanie Jegelka, and Ali Jadbabaie. 2025. On the Emergence of Position Bias in Transformers. *arXiv preprint arXiv:2502.01951* (2025).
- [67] Peng Xu, Wei Ping, Xianchao Wu, Lawrence McAfee, Chen Zhu, Zihan Liu, Sandeep Subramanian, Evelina Bakhturina, Mohammad Shoeibi, and Bryan Catanzaro. 2023. Retrieval meets long context large language models. In *The Twelfth International Conference on Learning Representations*.
- [68] Murong Yue, Jie Zhao, Min Zhang, Liang Du, and Ziyu Yao. 2024. Large Language Model Cascades with Mixture of Thought Representations for Cost-Efficient Reasoning. In *The Twelfth International Conference on Learning Representations (ICLR)*. <https://openreview.net/forum?id=6okaSfANzh>
- [69] Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in neural information processing systems* 36 (2023), 46595–46623.
- [70] Xuanhe Zhou, Guoliang Li, Jianming Wu, Jiesi Liu, Zhaoyan Sun, and Xinning Zhang. 2023. A learned query rewrite system. *Proceedings of the VLDB Endowment* 16, 12 (2023), 4110–4113.

Received July 2025; revised October 2025; accepted November 2025