

Example-Driven Intent Synthesis for Constrained Data Bundle Retrieval: Focused Text Snippet Extraction and Beyond

Whanhee Cho
University of Utah
whanhee.cho@utah.edu

Kuangfei Long
Boston University
longkuangfei@outlook.com

Mahmood Jasim
Louisiana State University
mjasim@lsu.edu

Matteo Brucato
OSM Data
matteo@osm-data.com

Alexandra Meliou
UMass Amherst
ameli@cs.umass.edu

Peter J. Haas
UMass Amherst
phaas@cs.umass.edu

Anna Fariha
University of Utah
afariha@cs.utah.edu

ABSTRACT

Selecting a *bundle* of items that collectively satisfies constraints is a fundamental task across databases, recommender systems, and text summarization. Unlike traditional top- k retrieval, bundle retrieval is inherently combinatorial and, in general, NP-hard. Although package queries can efficiently retrieve bundles given a well-formed query, two key user-centric challenges remain: (1) expressing and tuning multi-dimensional bundle intent through a user-friendly interface, and (2) ensuring feasibility when the query yields empty results. We introduce Ex2BUNDLE, an *Example-driven Bundle* retrieval framework that enables users to specify their intent through *example bundles* and automatically synthesizes package queries that capture the intent implicit in those examples via aggregate constraints. Ex2BUNDLE also addresses a challenge unique to bundle retrieval: when inferred constraints are infeasible over the target data, our *data-aware constraint relaxation* minimally adjusts the constraint bounds while preserving alignment with user intent. We instantiate a specific application of Ex2BUNDLE. Extensive experiments over real-world datasets and a user study demonstrate that Ex2BUNDLE improves usability and consistently returns intent-aligned bundles even under distributional shifts of the target database.

Artifact Link: <https://github.com/kuangfei-long/ex2bundle>.

1 INTRODUCTION

Data *bundle* retrieval is a fundamental task across multiple domains: package queries retrieve packages (sets of tuples) from relational databases under aggregate constraints [12]; recommender systems suggest playlists (sets of songs) or combo offers (sets of products) tailored to user preferences [5]; and extractive text summarization systems retrieve subsets of sentences from documents that form coherent summaries based on user queries [22, 66, 94]. Unlike traditional retrieval that returns individual or top- k items, bundle retrieval involves selecting a *set* (or package [12]) whose elements must collectively optimize global objectives while satisfying set-level constraints that align with user preferences. Consequently, the inclusion or exclusion of any single item affects the feasibility and overall quality of the bundle, making the problem inherently combinatorial and NP-hard [12]. While systems exist for the efficient retrieval of data bundles from this combinatorial search space [12], two user-centric challenges remain as primary obstacles: (1) the *interface* challenge—how can humans easily express their intent of the desired data bundles—and (2) the *feasibility* challenge—how to ensure non-empty answers for the user queries.

Challenge 1: Usable interface. The first challenge is communicating the user intent of the desired data bundle. Like SQL, the Package Query Language (PaQL) [12] allows specification of a bundle query through linear constraints and objective functions over data attributes, but the users must know the exact query parameters: attribute names and their numerical bounds. This is particularly difficult for non-experts who must (i) learn the PaQL [12] syntax, (ii) know the database schema, (iii) determine the attributes of interest, (iv) have a good sense of the value distributions, and (v) translate their intent precisely to parameterized constraints and global objectives of a PaQL query. We show in Example 1 that even for experts, step (v) is particularly difficult for PaQL because constraints are over *bundle aggregates* (SUM, AVG, COUNT): users must reason at the bundle level—predicting what a SUM should be over a bundle of unspecified size—rather than at the tuple level as in SQL. For certain applications, such as extractive document summarization [94], an alternative is to specify the intent in natural language. However, as prior work shows, natural language is too generic to express specific user intents [22] and subsequent conversational tuning of intent is imprecise and frustrating for humans [97].

EXAMPLE 1 (SYNTHESIZING A PACKAGE QUERY). *Morpheus, a CS department Chair, is seeking to fill two tenure-track assistant professor positions to strengthen AI and Databases. Since he plans for the two new hires to jointly teach an “AI in Databases” course, he wants them to have reasonable collective teaching experience, but not too much, as that may indicate teaching-focused or senior candidates who are unlikely to accept a tenure-track entry-level position. He also wants to maximize the total recommendation score of the new hires based on their letters. This is a bundle retrieval problem: Morpheus must pick a set of two out of 200 candidates that best satisfy his criteria, as shown via the following under-specified package query:*¹

```
Q1:  SELECT PACKAGE(*) FROM CANDIDATES
      SUCH THAT COUNT(*) = 2
      AND SUM(ai_score) is high
      AND SUM(db_score) is high
      AND SUM(teaching_score) is reasonable
      MAXIMIZE SUM(reco_score);
```

However, at this point, all Morpheus has are the raw application materials (resumes, statements, letters, etc.) of these candidates. Before querying the candidate database, he must first map each candidate’s application materials to a 4-dimensional vector space: ai_score,

¹While this toy example involves 2 hires over 4 criteria; faculty searches typically involve 5–10 candidates and 15+ criteria, yielding a combinatorial search space. E.g., choosing 5 from 200 candidates yields $\binom{200}{5} \approx 2.5$ billion possible bundles.

	AI ai_score	Databases db_score	Teaching teaching_score	Recommendation reco_score
Smith	0.8	0.4	0.1	0.4
Jones	0.4	0.7	0.2	0.9
Neo	0.4	0.5	0.4	0.8
Brown	0.5	0.4	0.4	0.9

Table 1: Partial list of candidates for Example 1.

db_score, teaching_score, and reco_score. Morpheus decides to use an off-the-shelf embedding technique to obtain numeric scores for each of these dimensions for every candidate (Table 1). However, he now faces another problem: he does not know what constitutes an appropriate replacement for the underspecified values *high* and *reasonable* in Q1. Is 0.7 considered a *high* score for AI expertise? What value represents the *reasonable* range for teaching score? He can examine the value distributions of the embeddings, but still cannot determine which corresponding values would accurately capture his intended criteria.

Example 1 highlights two key struggles in formulating package queries to express bundle query intent: (1) the lack of means to accurately map a database to a vector database with appropriate dimensions of interest, which primarily affects non-experts; and (2) the lack of knowledge to correctly parameterize the PaQL query, which affects experts and non-experts alike. In Example 1, Morpheus was interested in only 4 criteria; however, his struggle would be even worse if there were a larger number of criteria.

Our solution to Challenge 1: bundle-query by example. To overcome the usability challenge in the interface of bundle query intent specification, we build on the *Query by Example* (QbE) paradigm [101], which has seen significant success in traditional SQL querying [24] and other domains [22, 29, 31, 71] by lowering the barrier to task specification. We extend QbE to bundle queries (BQbE): the user provides a few exemplar data bundles to convey their query intents implicitly, thereby bypassing the construction of a precisely parameterized PaQL query. BQbE is particularly useful when users cannot articulate precise criteria but can instead provide representative examples of what they seek.

EXAMPLE 2 (PARAMETERIZING A PACKAGE QUERY). *Continuing from Example 1, Morpheus recalls that strong new hires resemble recent hires at top universities, such as {Trinity, Cypher} (University X) and {Link, Niobe, Seraph} (University Y). These exemplars excel in Databases and AI with reasonable collective teaching experience. Morpheus computes their scores across 4 dimensions using the same embedding of Example 1 (Table 2). Guided by the aggregated scores, he parameterizes the underspecified query Q1 to obtain:*

```
Q2: SELECT PACKAGE(*) FROM CANDIDATES
      SUCH THAT COUNT(*) = 2
      AND SUM(ai_score) BETWEEN 0.8 AND 1.2
      AND SUM(db_score) BETWEEN 1.0 AND 1.3
      AND SUM(teaching_score) BETWEEN 0.7 AND 1.3
      MAXIMIZE SUM(reco_score);
```

Example 2 highlights a scenario where a user cannot directly specify precise criteria via PaQL query parameters but can provide valid *examples* to implicitly convey the criteria. These examples can be used to automatically learn user preferences and translate them to query parameters. However, another key challenge remains, which we highlight in Example 3.

	AI ai_score	Databases db_score	Teaching teaching_score	Recommendation reco_score
Example 1: University X hires				
Trinity	0.6	0.5	0.3	0.7
Cypher	0.2	0.8	0.4	0.9
SUM	0.8	1.3	0.7	1.6

Example 2: University Y hires				
Link	0.9	0.3	0.4	0.6
Niobe	0.2	0.3	0.6	0.8
Seraph	0.1	0.4	0.3	0.7
SUM	1.2	1.0	1.3	2.1

Morpheus' (implicit) intent

[0.8, 1.2]	[1.0, 1.3]	[0.7, 1.3]	Maximize
------------	------------	------------	----------

Table 2: Example bundles help discover query parameters.

bundle	ai	db	teaching	reco	valid?
b ₁ {Smith, Jones}	1.2	1.1	0.3 (↓ 0.4)	1.3	no
b ₂ {Smith, Neo}	1.2	0.9 (↓ 0.1)	0.5 (↓ 0.2)	1.2	no
b ₃ {Smith, Brown}	1.3 (↑ 0.1)	0.8 (↓ 0.2)	0.5 (↓ 0.2)	1.3	no
b ₄ {Jones, Neo}	0.8	1.2	0.6 (↓ 0.1)	1.7	almost
b ₅ {Jones, Brown}	0.9	1.1	0.6 (↓ 0.1)	1.8	almost
b ₆ {Neo, Brown}	0.9	0.9 (↓ 0.1)	0.8	1.7	almost

Table 3: No candidate bundle fully satisfies the constraints of Q2. For instance, b₆ violates the constraint “SUM(db_score) BETWEEN 1.0 AND 1.3”, as its db_score (0.9) is 0.1 below the lower bound (1.0).

EXAMPLE 3 (ENSURING QUERY FEASIBILITY TO OBTAIN RESULTS). *Morpheus issues Q2 to a package query execution engine [12], but it returns no result. A close inspection reveals that indeed none of the 2-candidate bundles fully satisfies Q2 (Table 3). However, the last three bundles “almost” satisfy the constraints. If the db_score constraint had a slightly relaxed lower bound (0.9 instead of 1.0), then b₆ would be a valid bundle. Similarly, slightly relaxing the lower bound of the teaching_score constraint (0.6 instead of 0.7) would result in two additional valid bundles: b₄ and b₅. Then following the goal of maximizing the reco_score, the bundle b₅ would be the final result.*

Challenge 2: ensuring feasibility. While BQbE addresses the first challenge, it brings forth a second one that is particular to bundle retrieval and absent from traditional single-tuple QbE: constraints inferred from example bundles may yield queries that are *infeasible* over the target data, returning empty answers [56, 65], because aggregate constraints can be jointly unsatisfiable even when individually plausible, especially under distribution shift between the example and target domains. In Example 3, even with a fully formed PaQL query, with the desired constraints and specific parameters, it turned out to be *infeasible* over the target data (Morpheus’ University). This happened because the data distributions differed between the example domains (Universities X and Y) and the target domain (Morpheus’ University), a common scenario in practice [42, 74]. Existing package query execution systems [12] provide no further insight into how the query parameters interact with the target data, leaving users uncertain about how to tune the parameters to make the query feasible such that valid results are retrieved.

Our solution to Challenge 2: data-aware relaxation of query parameters. To ensure feasibility even when the original query is infeasible, we introduce *data-aware constraint relaxation techniques* that adapt the synthesized PaQL query to the target data

by adjusting constraint bounds. By analyzing the data distribution of the target database, we identify which constraints act as bottlenecks that prevent “almost-valid” bundles from being retrieved. This analysis guides us to relax the appropriate constraints, while minimizing deviation from the user’s original intent.

Why existing approaches fall short. Three classes of relevant approaches exist, but all of them fail to support BQbE.

PaQL and QbE systems. PaQL engines [12] require fully parameterized queries and return empty results when infeasible; QbE systems [24, 101] reduce the parameter burden for single-tuple retrieval but do not extend to bundle queries with aggregate constraints.

Top- k retrieval. For BQbE, an alternative is to treat the entire example bundle as a single query vector and retrieve top- k similar tuples based on vector similarity [55, 78] to assemble the result bundle. However, this approach can miss globally optimal bundles and violate intended constraints, since top- k retrieval evaluates items independently using an implicit scoring function—without guarantees of global optimality—whereas bundle retrieval requires reasoning over combinations of tuples under globally enforced [100], multi-dimensional aggregate constraints. RAG systems inherit the same limitation, as they perform top- k nearest-neighbor search over individual tuples (but not their combinations) before passing results to a generative model. Another alternative is to match each tuple in the example bundle independently and combine their top matches to form the result bundle; yet this approach breaks down when the example and target bundle sizes differ. In summary, bundle retrieval is fundamentally an NP-hard combinatorial optimization problem, and similarity-based greedy top- k search fails to address it.

Conversational LLMs. An LLM prompted with examples (few-shot learning) could produce the bundle directly, but this requires reasoning over the entire target data and solving a combinatorial optimization problem, which LLMs handle unreliably [88]. An LLM could instead synthesize PaQL queries; however, the bounds would reflect an opaque mapping between the user’s examples and constraint values. LLMs also have no native mechanism to detect infeasibility or relax bounds in an intent-preserving way, and per-query latency limits interactive use. In a pure NL interface, without examples to anchor the bounds, the user must rely on conversational refinement, which is imprecise and tedious [97].

Desiderata. Examples 1–3 and the limitations of alternative approaches motivate the desiderata of an ideal BQbE system:

- D1** Ensure a user-friendly interface for specifying intents, where users can provide examples to implicitly convey their bundle query intents, without requiring any knowledge of the database schema and complex query syntax and parameters.
- D2** Model the user intent implicit in the examples into a representation with explicit parameters, such as a PaQL query.
- D3** Ensure feasibility of the synthesized PaQL query over the target database, yielding a non-empty result even when no bundle perfectly matches the user’s intent.
- D4** Provide an intuitive interface for iterative intent refinement, allowing users to adjust query constraints.

Ex2BUNDLE. We introduce Ex2BUNDLE, an example-driven data bundle retrieval system that (i) enables users to specify bundle query

intent through example bundles, addressing D1; (ii) synthesizes a package query from these examples to transparently capture user intent, addressing D2; (iii) applies data-aware constraint relaxation to ensure feasibility of the synthesized query w.r.t. the target data domain, addressing D3; and (iv) provides an intuitive interface for interactive intent refinement, addressing D4.

Use cases. We now present three real-world applications where BQbE lowers the barrier for data bundle retrieval:

- *Supplier selection for business.* Businesses entering new markets struggle to identify a set of suppliers via explicit constraints over acceptable price ranges, inventory availability, or financial stability thresholds. Ex2BUNDLE can infer these constraints implicitly from examples of supplier bundles drawn from existing markets where the business is successful, adapt them to the target market, and retrieve an optimal set of suppliers that satisfies the constraints.
- *Focused snippet extraction.* In the context of text summarization, BQbE arises as *focused text snippet extraction*: given a “focus” in the form of several source document–summary pairs, where each summary is expressed as a set of extracted sentences (snippet), the goal is to select a snippet from a target document that is consistent with the examples. Ex2BUNDLE builds on our prior work SuDocu [22] to support personalized extractive summarization [95].
- *Playlist recommendation.* Existing playlist recommenders—e.g., Spotify’s Discover Weekly [80], YouTube’s Daily Discover [8], and Netflix’s Top Picks [3]—rely on item-level or single-playlist similarity, which often fail to capture diverse tastes and cause user frustration [1]. Ex2BUNDLE allows users to provide example playlists reflecting their desired spread across genres, artists, tempo, etc., and generates playlists that satisfy their preferences.

Contributions. We make the following contributions:

- We formalize the problem of *example-driven bundle retrieval* and show how it translates to the package query framework (§2).
- We present Ex2BUNDLE, an end-to-end framework for *example-driven bundle retrieval* that synthesizes PaQL query constraints from user examples. We contribute a *data-aware constraint-bound relaxation technique* that ensures query feasibility. We further introduce design principles for an interactive slider-based interface for intent refinement, along with effective techniques to translate between sliders and constraint bounds (§3).
- We evaluate Ex2BUNDLE on two use cases—package queries over the TPC-H dataset [87] and focused text snippet extraction over real-world datasets [34, 95]—showing that Ex2BUNDLE achieves 100% constraint satisfaction while maintaining competitive objective scores, and outperforms retrieval and extractive baselines. Notably, Ex2BUNDLE remains competitive with LLM-based approaches without incurring any additional token or inference cost, while scaling to realistic workloads (§4).
- Our user study shows that Ex2BUNDLE achieves high user satisfaction, due to improved ease of use and customizable sliders, for the task of focused text snippet extraction (§5).

2 PROBLEM FORMULATION

In this section, we formalize the problem of example-driven bundle retrieval. We begin by introducing how we model the *source* and *target data*—from which bundles are retrieved—and *example*

Symbol	Description
$\mathbf{f} = \{f_1, \dots, f_K\}$	Common schema over K attributes/features
$T_i^s \models \mathbf{f}$	A source data over the schema $\mathbf{f}$
$t \models \mathbf{f}$	A single tuple over the schema $\mathbf{f}$
$\mathbf{f}(t) = [f_1(t), \dots, f_K(t)] \in \mathbb{R}^K$	Feature vector of tuple t
$E_i \subseteq T_i^s$	Example bundle for the source data T_i^s
$\mathbf{T}^s = \{T_1^s, \dots, T_N^s\}$	Set of N source data
$\mathbf{E} = \{E_1, \dots, E_N\}$	Set of N user-provided example bundles
$T^q \models \mathbf{f}$	A target data over $\mathbf{f}$ to retrieve a bundle from
$B \models \mathbf{f}$	A bundle over the schema $\mathbf{f}$
$\Theta = \{\langle \text{lb}_j, \text{ub}_j \rangle\}_{j=1}^K$	Constraint bounds for K attributes
$\mathbf{F}(B) = [F_1(B), \dots, F_K(B)]$	Feature profile of bundle B
$\mathbf{F}(B) \vdash \Theta$	$\mathbf{F}(B)$ satisfies the constraints given by Θ
$\sigma(t) \mapsto \mathbb{R}$	Domain-specific tuple-level scoring function

Table 4: Table of notations. Bold letters denote sets or vectors.

bundles—which communicate the user’s intent. We then show how bundle retrieval can be naturally modeled within the package query framework and define the problem of *example-driven package query synthesis* for bundle retrieval, our focus in this paper. Table 4 provides a summary of notations used throughout the paper.

Source and target data. In example-driven bundle retrieval, each example bundle is taken from a corresponding source data. In Example 2, {Trinity, Cypher} is an example bundle from the source data of University X and {Link, Niobe, Seraph} from the source data of University Y. One key requirement here is that the source data for all the example bundles must have the same schema, to allow for effective intent discovery across the shared attributes.

We define a schema $\mathbf{f}$ of a single table² over K numerical³ attributes/features $\{f_1, \dots, f_K\}$. Each source data T_i^s must conform to $\mathbf{f}$, denoted as $T_i^s \models \mathbf{f}$. In Table 2, the source data for all the Universities conform to the schema {ai_score, db_score, teaching_score, reco_score}. A tuple t over $\mathbf{f}$, denoted as $t \models \mathbf{f}$, is characterized by a feature vector $\mathbf{f}(t) = [f_1(t), \dots, f_K(t)] \in \mathbb{R}^K$. In Table 2, Trinity’s feature vector is [0.6, 0.5, 0.3, 0.7]. We use $T^q \models \mathbf{f}$ to denote the *target data* from which the user wants to extract their desired bundle. In Example 2, Table 1 represents the target data.

Example bundles. In example-driven bundle retrieval, users convey their intent via a set of *example bundles* $\mathbf{E} = \{E_1, \dots, E_N\}$, where each $E_i \subseteq T_i^s$ consists of a subset of tuples from the corresponding source data T_i^s . In Example 2, $\mathbf{E} = \{\{\text{Trinity, Cypher}\}, \{\text{Link, Niobe, Seraph}\}\}$ and the source data set $\mathbf{T}^s = \{T_{\text{UoF}}^s, T_{\text{UoY}}^s\}$.

2.1 Example-driven bundle retrieval

We are now ready to (informally) define our problem:

PROBLEM 2.1 (EXAMPLE-DRIVEN BUNDLE RETRIEVAL). *Given a set of source data $\mathbf{T}^s$ over the same schema $\mathbf{f}$, corresponding example bundles $\mathbf{E}$, where $E_i \in \mathbf{E}$ is an example bundle for the source data $T_i^s \in \mathbf{T}^s$, and a single target data $T^q \models \mathbf{f}$, identify a bundle $B^* \subseteq T^q$ that (i) satisfies the “user intent” implicit in $\mathbf{E}$ w.r.t $\mathbf{T}^s$ and (ii) is the “best” among such bundles.*

However, this problem is underspecified: it is unclear how to model the user intent and define “best” or the notion of optimality.

²While our formalization focuses on a single-table schema, it is not an inherent limitation: we support relational databases by joining tables into a denormalized table, as shown in our experiments over the TPC-H dataset (§4).

³For unstructured data such as text, we derive numerical features through domain-specific methods such as topic modeling or learned vector embeddings (§3) to produce a structured representation under a fixed numerical schema.

2.1.1 Modeling user intent. Capturing user intent from example bundles may require complex models involving non-linear constraints. However, practical use cases share the following common structure: (i) the desired bundle is a *set* of items, (ii) quality of the bundle depends on some *aggregate* properties (e.g., total score) of the set, and (iii) the user preference is expressible as *bounds* on those aggregates. These align perfectly with *package query* [12]—which retrieves subsets of tuples from relational tables that collectively satisfy (linear) aggregate constraints while optimizing a (linear) global objective—as illustrated by Q2 in Example 2. Thus, we model user intent as *linear constraints* on the bundle’s *feature profile*, with an optional cardinality constraint on the bundle size. The benefit of this modeling is twofold: it provides *interpretability*—as linear constraints can be easily understood and inspected by humans [32, 62, 63, 69, 73, 75]—and *tractability*—it maps to Integer Linear Programming (ILP) with established solvers and foundations [12].

Feature profile of a bundle. To capture the feature-wise properties of a bundle B , we define its *feature profile* $\mathbf{F}$ as a vector obtained by aggregating, for each feature, the values of that feature across all tuples in B . Formally,

$$\mathbf{F}(B) = [F_1(B), \dots, F_K(B)], \text{ where } F_j(B) = \mathcal{A}_{t \in B} F_j(t)$$

Here, $\mathcal{A}$ is an aggregate such as SUM. In Table 2, using SUM as the aggregate, the feature profile of the bundle {Trinity, Cypher} is [0.8, 1.3, 0.7, 1.6]—the column-wise sums for University X hires.

Formulating constraints. Following prior work in QbE [24], given a set of user-provided example bundles, we posit that the user’s intended result bundle would exhibit a feature profile similar to that of the example bundles. Accordingly, we model the user’s intent as a *set of bounded constraints over the feature profile* of the desired bundle, where the derivation of the bounds should be guided by the feature profiles of the example bundles in $\mathbf{E}$, the source data set $\mathbf{T}^s$, and the target data T^q . Formally:

$$\{ \text{lb}_1 \leq F_1(B) \leq \text{ub}_1, \dots, \text{lb}_K \leq F_K(B) \leq \text{ub}_K \}$$

where lb_j and ub_j depend on $\mathbf{E}$, $\mathbf{T}^s$, and T^q for $1 \leq j \leq K$

We use $\Theta = \{\langle \text{lb}_1, \text{ub}_1 \rangle, \dots, \langle \text{lb}_K, \text{ub}_K \rangle\}$ to denote constraint bounds for all K attributes over the schema $\mathbf{f}$. The notation $\mathbf{F}(B) \vdash \Theta$ denotes that the feature profile of the bundle B satisfies the constraint bounds specified by Θ , i.e., $\forall 1 \leq j \leq K, \text{lb}_j \leq F_j(B) \leq \text{ub}_j$.

2.1.2 Defining bundle optimality. The second issue in Problem 2.1 is that, when multiple candidate bundles satisfy Θ , a principled way is required to quantify their quality and ensure optimality. To this end, we assume knowledge of an application-specific tuple-level function $\sigma(t) \mapsto \mathbb{R}$, typically provided by a domain expert who configures the system for end users. We define the quality of a bundle B by aggregating σ over all $t \in B$, which naturally fits the linear objective function optimized by the PaQL framework. For example, this could correspond to maximizing the total recommendation score in Example 1, or minimizing the word count in text summarization. Alternatively, σ can be learned from user feedback on bundle quality, e.g., for playlist recommendation in music streaming.

2.2 Example-driven package query synthesis

With our models of user intent and bundle optimality, which naturally align with the package query framework, we reformulate

Problem 2.1 of example-driven bundle retrieval as that of synthesizing a package query—specifically, its parameters—from example bundles. The synthesized package query serves as a *mechanism* for efficiently retrieving the optimal result bundle.

Given a target data T^q , constraint bounds $\Theta = \{\langle lb_j, ub_j \rangle\}_{j=1}^K$, and (optional) cardinality constraint bounds $C = \langle lb_c, ub_c \rangle$, we fix the following parameterized package query:

```
PQ( $T^q, \Theta, C$ ): SELECT PACKAGE(*) AS B FROM  $T^q$ 
SUCH THAT COUNT(B) BETWEEN  $lb_c$  AND  $ub_c$ 
AND  $F_j(B)$  BETWEEN  $lb_j$  AND  $ub_j \forall 1 \leq j \leq K$ 
MAXIMIZE  $\mathcal{A} \sigma(t)$ 
 $t \in B$ 
```

Here, $\mathcal{A}$ aggregates the tuples $t \in B$ using σ to compute the quality of B , which the package query aims to maximize as its objective. Without loss of generality, a minimization objective can be expressed as a maximization objective by simply negating the objective. The choice of aggregation functions (including those used to compute the feature profile $F(B)$) and the scoring function σ is typically guided by the application domain. These are system parameters determined during setup and are not optimized over. We discuss how to choose the aggregation and scoring functions in §3.3. Also, without loss of generality, the cardinality constraint $lb_c \leq \text{COUNT}(B) \leq ub_c$ can be subsumed into the feature-profile constraints by augmenting the feature profile with an additional dimension representing bundle cardinality. Hence, we do not explicitly include the cardinality constraint in the remainder of the paper.

In Example 2, Q2 is a package query over the schema $f = \{\text{ai_score}, \text{db_score}, \text{teaching_score}\}$, so $K = 3$. The aggregation used to compute the feature profile is SUM, and the objective is to maximize SUM(reco_score) (i.e., $\sigma = \text{reco_score}$).

While Morpheus manually derived the constraint bounds in Θ , the resulting query was infeasible over the target data T^q representing his university. Thus, the main requirement is to determine the parameter Θ , and use it to synthesize a package query that will retrieve the optimal bundle $B \subseteq T^q$, as specified in Problem 2.1.

PROBLEM 2.2 (EXAMPLE-DRIVEN PACKAGE QUERY SYNTHESIS). *Given a schema f , a set of source data $T^s = \{T_1^s, \dots, T_N^s\}$ where each $T_i^s \models f$, corresponding example bundles E , a target data $T^q \models f$, and a tuple-level scoring function $\sigma : t \mapsto \mathbb{R}$, together with the corresponding aggregator $\mathcal{A}$, determine a synthesis function $G : (E, T^s, T^q) \mapsto \langle \mathbb{R}, \mathbb{R} \rangle^K$ to derive the constraint bounds Θ such that:*

- (1) **Feasibility:** the package query $PQ(T^q, \Theta)$ returns a non-empty result over T^q , ensuring $\exists B \subseteq T^q$ s.t. $F(B) \vdash \Theta$.
- (2) **Optimality:** the retrieved bundle $B^* = PQ(T^q, \Theta)(T^q)$ maximizes the objective, i.e., $B^* = \arg \max_{B \subseteq T^q \text{ s.t. } F(B) \vdash \Theta} \mathcal{A} \sigma(t)$.

The core challenge here is to determine the synthesis function $G : (E, T^s, T^q) \mapsto \langle \mathbb{R}, \mathbb{R} \rangle^K$ to derive the constraint bounds Θ by capturing the intent implicit in E while generalizing to the target data T^q . Unlike general query-by-example, we do not focus on synthesizing package queries with arbitrary structures, which would require inferring joins and aggregates. Instead, we focus on inferring the parameters of the constraint bounds Θ . This task is particularly challenging because the distribution of T^q may differ from T^s : overly tight constraints (Example 2) can yield infeasible queries on the target data, while overly relaxed constraints allow bundles that deviate from user feasibility.

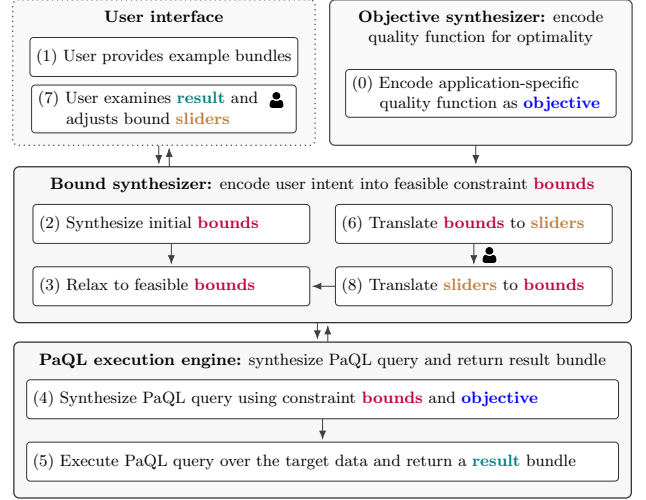


Figure 5: Ex2BUNDLE workflow: before end-user interaction, a domain expert provides a quality function to be used as the objective and specifies an aggregate function to be used in the constraints. (0) Ex2BUNDLE encodes the quality function as objective, (1) the user provides example bundles, from which Ex2BUNDLE formulates constraints using the expert-specified aggregate, (2) synthesizes initial constraint bounds and (3) relaxes them to ensure feasibility. A PaQL query is (4) formed using these bounds and the objective and (5) executed to retrieve a result bundle. For further intent refinement, (6) an interactive slider interface is generated, (7) the user inspects results and adjusts sliders, (8) Ex2BUNDLE maps adjustments back to bounds, and repeats (3)–(5).

3 THE EX2BUNDLE FRAMEWORK

Ex2BUNDLE consists of four main components (Fig. 5): (§3.1) *bound slider*, an interface that allows users to refine their intent by adjusting constraint bounds; (§3.2) *bound synthesizer*, which encodes user intent into *feasible* constraint bounds; (§3.3) *objective synthesizer*, which encodes an application-specific quality function as the objective of the PaQL query; and (§3.4) *PaQL execution engine*, which synthesizes and executes package queries to retrieve result bundles.

Running example for focused text snippet extraction. Throughout the rest of the paper, we use a running example centered on the task of *Focused Text Snippet Extraction* (FTSE), a primary application through which we evaluate Ex2BUNDLE (§4). Because Ex2BUNDLE requires numeric features and text is non-numeric, we apply *contextualized topic modeling* (CTM) [6] with sentence embeddings from SBERT [84] to map sentences into a numeric vector space of topic-based features. Under the resulting topic model defined over schema f , $f_j(t) \in [0, 1]$ denotes the relevance of a sentence t to topic f_j , and $f(t)$ denotes the vector representation of t in the topic space f . For ease of understanding, we use the following, more specific and domain-relevant terms for FTSE, instead of generic terms:

General term	FTSE term	General term	FTSE term
source/target data	source/target document	tuple	sentence
bundle	snippet or summary	schema	topic model
attribute/feature	topic	feature profile	topic profile

EXAMPLE 4. *Consider a topic model f over Demographics, Geography, and Economy applied to documents describing U.S. states. Let T_{UT}^s denote the Utah document and $E_{UT} = \{t_{UT}^1, t_{UT}^2, t_{UT}^3\} \subseteq T_{UT}^s$ be three*

Sentence/snippet	Topic score/profile
(t_{UT}^1) The largest ancestry groups in the state are: 26.0% English, 11.9% German, [...]	
(t_{UT}^2) In comparison to all the U.S. states and territories, Utah, with a population of just over three million, is the 13th-largest by area, the 30th-most populous, and the 11th-least densely populated.	
(t_{UT}^3) St. George was the fastest-growing metropolitan area in the United States from 2000 to 2005.	
(E_{UT}) The largest ancestry groups in the state are: 26.0% English, 11.9% German, [...]. In comparison to all the U.S. states and territories, Utah, with a population of just over three million, is the 13th-largest by area, the 30th-most populous, and the 11th-least densely populated. St. George was the fastest-growing metropolitan area in the United States from 2000 to 2005.	

Table 6: Topic scores for sentences t_{UT}^1 , t_{UT}^2 , and t_{UT}^3 ; and topic profile for the example snippet $E_{UT} = \{t_{UT}^1, t_{UT}^2, t_{UT}^3\}$.

sentences (Table 6). Under $\mathbf{f}$, t_{UT}^1 maps to $\mathbf{f}(t_{UT}^1) = [0.90, 0.00, 0.00]$, indicating a *demographics* focus; t_{UT}^2 to $[0.90, 0.90, 0.00]$, covering *demographics* and *geography*; and t_{UT}^3 to $[0.90, 0.30, 0.10]$, primarily *demographics* with weaker *geography* and *economy* associations.

3.1 Bound slider interface for intent refinement

A key requirement (D4 in §1) for Ex2BUNDLE is an intuitive user interface for iterative intent refinement. A simple approach is to let users directly edit the synthesized constraint bounds (Ex2BUNDLE’s bound synthesis is discussed in §3.2). However, non-experts often do not know the precise mapping between their intent and constraint bounds (as in Example 1). Based on studies [16, 25, 44, 76, 81, 92] establishing that humans are better at relative judgments (more vs. less) than absolute ones (assigning ratings) [86], we introduce a single-parameter slider that abstracts the raw bounds.

Fig. 7 shows the slider interface, where each position corresponds to a pair of upper and lower bounds, and the neutral point corresponds to the feasible bounds synthesized from user examples (details in §3.2). The slider instance in Fig. 7 depicts the topic DEMOGRAPHICS, where the neutral point 0 corresponds to the feasible bounds $\langle 1.50, 2.70 \rangle$. The slider further encodes other representations on a scale from -100 to $+100$ around the neutral point to allow fine-grained adjustments. Upon reviewing the result bundle, users simply indicate whether they want more (less) relevance to the topic by moving the slider right (left). This design also reduces the number of controllable parameters, consistent with HCI findings that simpler control spaces are cognitively less demanding [39].

Mapping between slider positions and bounds. We use γ_j to denote the position of the slider for topic f_j . We map the initial feasible bounds $\langle lb_j, ub_j \rangle$ —synthesized from user examples—to the midpoint of the slider at $\gamma_j=0$. For a target document T^q , the minimum and maximum attainable bundle scores for topic f_j are 0 and $F_j(T^q)$, respectively. Let us use mx to denote $F_j(T^q)$ for simplicity. These two extreme values—0 and mx —map to the lower bound of the left-most slider position and upper bound of the right-most slider position. A key challenge here is that both sides of the midpoint of the slider have an equal number of positions (100 on each side); however, the distance between the initial feasible bounds and the two extreme bounds is not necessarily equal. Thus, for slider positions

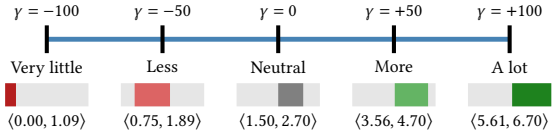


Figure 7: Slider for DEMOGRAPHICS. Users can adjust topic emphasis from -100 (very little) to $+100$ (a lot). Each slider position maps to specific constraint bounds; the neutral position (0) corresponds to the feasible bounds synthesized (and relaxed) from user examples.

left of the midpoint, we determine the *lower* bound by interpolating based on the position’s relative distance between the left-most endpoint ($\gamma_j=-100$) and the midpoint ($\gamma_j=0$). Symmetrically, for positions right of the midpoint, we determine the *upper* bound by interpolating based on the position’s relative distance between the right-most endpoint ($\gamma_j=+100$) and the midpoint ($\gamma_j=0$). Assuming a constant width w for all bounds, the mappings are as follows:

Slider Position	Lower bound	Upper bound	γ_j
(left-most)	0	w	-100
(left of the center)	$lb_j \cdot \frac{(100-x)}{100}$	$lb_j \cdot \frac{(100-x)}{100} + w$	$-x$
(center)	lb_j	ub_j	0
(right of the center)	$ub_j + \frac{(mx-ub_j) \cdot x}{100} - w$	$ub_j + \frac{(mx-ub_j) \cdot x}{100}$	$+x$
(right-most)	$mx - w$	mx	$+100$

A question still remains: how to determine the value of w ? Should it be constant for all slider positions? We use the width of the neutral position ($ub_j - lb_j$) as a guide to derive the bound width of slider position at $\gamma_j=x$ with the formula: $w(x) = e^{-\frac{\alpha \cdot |x|}{100}} \cdot (ub_j - lb_j)$.

Here, α is a tunable parameter between 0 and ∞ , which determines how aggressively bound widths are narrowed as the slider moves farther away from the neutral position. When $\alpha=0$, all slider positions use the same width as the neutral position, i.e., $(ub_j - lb_j)$. As α increases, the width shrinks more rapidly as the slider moves away from the neutral position, resulting in increasingly tighter bounds for extreme values of γ_j . Regardless, Ex2BUNDLE automatically “snaps” the slider to a feasible position in case the user’s adjustment leaves it at an infeasible state. Note that the bounds represented by different slider positions can be overlapping (Fig. 7).

Overall, the slider interface offers the following benefits: (1) *simplicity* by reducing the number of controls, (2) *semantic clarity* through an intuitive scale from less (negative) to more (positive), (3) *feasibility preservation* by automatically snapping user adjustments to a feasible state, and (4) *reversibility* by allowing users to reset to the original (or any intermediate) intent. Users can still adjust the bounds directly if they wish, but the slider provides a more intuitive alternative that bypasses the complex underlying numeric bounds. We evaluate the effectiveness of the slider interface in §5.

3.2 Encoding intent into feasible constraints

The bound synthesizer (i) synthesizes initial, potentially infeasible, bounds from user examples to form constraints, (ii) relaxes the initial bounds to ensure feasibility, (iii) exposes the feasible bounds via the slider interface for iterative refinement, and (iv) translates the adjusted sliders back to feasible bounds after refinement.

3.2.1 Synthesizing the initial bounds. We extend SQUID [24], which learns user intent from example tuples, to a bundle-level granularity. Given example bundles $\mathbf{E} = \{E_i\}_{i=1}^N$, we use SUM as the aggregate to

Algorithm 1: Ex2BUNDLE bound relaxation algorithm

Input : Initial bounds Θ_{init} , schema f , target data T^q
Relaxation parameters μ and τ
Output : Feasible bounds Θ to construct the PaQL constraints

```

1  $\Theta \leftarrow \Theta_{\text{init}}$ 
2 #attempts  $\leftarrow 0$  // Initialize #attempts
3 while  $PQ(\Theta, T^q) = \emptyset$  // While query is infeasible
4 do
5    $\Theta_{\text{violated}} = \text{IdentifyViolations}(\Theta, T^q)$  // Violated constraints
6    $\Theta_{\text{relaxed}} = \emptyset$  // Relaxed constraints
7    $\rho \leftarrow \exp\left(\mu \cdot \max\left(1, \left\lfloor \frac{\text{\#attempts}++}{\tau} \right\rfloor\right)\right)$  // Relaxation step multiplier
8   foreach  $(lb_j, ub_j) \in \Theta_{\text{violated}}$  do
9      $\epsilon_j \leftarrow \frac{F_j(T^q)}{|T^q|}$  // Relaxation step size
10     $lb_j \leftarrow \max(lb_j - \rho \cdot \epsilon_j, 0)$  // Decrease lower bound
11     $ub_j \leftarrow \min(ub_j + \rho \cdot \epsilon_j, F_j(T^q))$  // Increase upper bound
12     $\Theta_{\text{relaxed}} = \Theta_{\text{relaxed}} \cup \{(lb_j, ub_j)\}$ 
13    $\Theta = (\Theta \setminus \Theta_{\text{violated}}) \cup \Theta_{\text{relaxed}}$  // Update constraint bounds
14 return  $\Theta$ 

```

compute their feature profiles $\{F(E_i)\}_{i=1}^N$. The choice of aggregate is application-specific and our framework can support any linear aggregate such as COUNT, MIN, MAX, etc. However, SUM is often preferred since AVG cannot distinguish single- from multi-tuple bundles, and MAX/MIN are outlier-sensitive. We derive the initial constraint bounds $\Theta_{\text{init}} = \{(lb_j, ub_j)\}_{j=1}^K$ by taking the topic-wise minimum and maximum of the feature profiles of bundles in E :

$$lb_j = \min_{1 \leq i \leq N} F_j(E_i), \quad ub_j = \max_{1 \leq i \leq N} F_j(E_i)$$

EXAMPLE 5. Consider 3 example snippets E_{UT} , E_{AZ} , and E_{WA} over the source documents T_{UT}^s , T_{AZ}^s , and T_{WA}^s that represent Wikipedia pages for Utah, Arizona, and Washington, respectively. The topic profile of E_{UT} , $F(E_{UT}) = f(t_{UT}^1) + f(t_{UT}^2) + f(t_{UT}^3) = [2.70, 1.20, 0.10]$, as shown in the last row of Table 6. The topic profiles for the example snippets, along with the initial constraint bounds are given below, which results in $\Theta_{\text{init}} = \{(1.50, 2.70), (0.90, 1.20), (0.10, 0.40)\}$.

	F_1 (DEMOGRAPHICS)	F_2 (GEOGRAPHY)	F_3 (ECONOMY)
$F(E_{UT})$	2.70 $\uparrow$	1.20 $\uparrow$	0.10 $\downarrow$
$F(E_{AZ})$	1.80	0.90 $\downarrow$	0.40 $\uparrow$
$F(E_{WA})$	1.50 $\downarrow$	1.00	0.30
lb (min)	1.50	0.90	0.10
ub (max)	2.70	1.20	0.40

Remark on constraint expressivity. Ex2BUNDLE supports any tuple-level constraint that standard SQL supports, such as equality predicates involving categorical attributes, logical conditions, and constraints over derived (potentially non-linear) attributes. Techniques for deriving such tuple-level constraints are addressed in prior work [24]. Ex2BUNDLE also supports more complex bundle-level constraints, such as interactions among aggregated attributes, and special cases of bounded constraints, such as equalities (by setting $lb = ub$) and one-sided inequalities (by setting $lb = -\infty$ or $ub = \infty$). Ex2BUNDLE supports any linear aggregate within the bundle-level constraints. The choice of the aggregate is application-specific and should be specified by a domain expert. In the absence of domain expertise, an LLM can suggest a suitable aggregate based on the task context. We note that practical applications may involve non-linear constraints, such as Max-Min diversity [4]. However, this is a limitation of PaQL, which Ex2BUNDLE relies on, not of Ex2BUNDLE itself.

In fact, the underlying ILP solvers such as IBM CPLEX [41] support certain classes of quadratic constraints with specific properties (e.g., convexity and semi-definiteness) and Ex2BUNDLE can potentially support such constraints once PaQL provides support for them.

3.2.2 Bound relaxation to ensure feasibility. Overly tight bounds can yield infeasible queries that return no valid bundle (Example 3), particularly when bounds are misaligned with the target data’s feature distribution under distributional shift (Example 3). Infeasibility arises from two sources. The first is *insufficient feature coverage*: the target data cannot satisfy a feature’s lower bound even when all tuples are selected. For example, if the source data (e.g., Utah, Arizona, & Washington) and example snippets emphasize national parks (e.g., Bryce, Grand Canyon, Olympics), but the target document is Kansas, which contains no national parks, its maximum contribution to that topic is 0, making any positive lower bound (e.g., 0.10) infeasible. The second is *atomicity*: tuples are indivisible, so selecting any single tuple may violate an upper bound. For instance, if every target sentence has a topic score of at least 0.15, an upper bound of 0.10 is infeasible because no tuple can be partially selected.

Bound relaxation algorithm. We address infeasibility via *bound relaxation*, which iteratively widens constraint bounds until the resulting package query becomes feasible. Algorithm 1 outlines this process, which repeats until feasibility is achieved (lines 5–13). If the initial constraint bounds Θ_{init} are infeasible (line 3), we first identify the violated constraints Θ_{violated} (line 5) using ILP solver signals (e.g., IBM CPLEX [41]). To preserve the original intent, we relax only the violated constraints. Specifically, we decrease lb_j toward 0 (line 10) and increase ub_j toward $F_j(T^q)$ (line 11). We determine the symmetric relaxation amount at each iteration using $\rho \cdot \epsilon_j$ (lines 10 & 11), where the step multiplier ρ (line 7) controls the relaxation rate based on the number of prior attempts, and the step size ϵ_j (line 9) is a target-data-specific parameter. We also considered two alternative relaxation techniques, including directional relaxation, which leverages the ILP solver’s information about whether the lower or upper bound causes a violation. However, as shown in §4.6.2, these alternatives never improved bundle quality over symmetric relaxation and had comparable runtime for typical relaxation severity. We therefore adopt symmetric relaxation due to its simplicity and empirical effectiveness.

Relaxation step multiplier, ρ . It controls the relaxation rate and is defined as $e^{\mu \cdot \max(1, \lfloor \text{\#attempts}/\tau \rfloor)}$. Here, μ (default 0.5) sets the base relaxation rate, while τ (default 10) controls how frequently ρ is boosted. Consequently, ρ increases exponentially every τ attempts, enabling faster escape from infeasible regions.

Relaxation step size, ϵ_j . While ρ determines the rate of relaxation, different features exhibit varying score distributions within T^q . To account for this, we derive a feature-specific step size from the target data distribution and set $\epsilon_j = \frac{F_j(T^q)}{|T^q|}$, which corresponds to the average contribution of feature f_j per tuple in T^q . Since the output bundle is a subset of the target data T^q , the appropriate relaxation magnitude depends on the empirical scale of each feature in T^q . We therefore derive ϵ_j from T^q : a fixed step can be too small or too large for features with different value ranges. Computing ϵ_j per feature automatically adapts the step size to the scale of that feature, even when feature values are not normalized.

EXAMPLE 6. Consider the California document T_{CA}^q over 5 sentences ($|T_{CA}^q| = 5$), where $F_2(T_{CA}^q) = 0.60$. The initially synthesized constraint is $0.90 \leq F_2(B) \leq 1.20$ (Example 5). This is infeasible since the maximum achievable score over T_{CA}^q is 0.60, requiring relaxation. For the first 10 iterations, $\rho = e^{0.5 \times 1} \approx 1.65$ and $\epsilon = \frac{0.60}{5} = 0.12$. Thus, the first iteration relaxes the bounds to $\max(0.90 - 1.65 \times 0.12, 0) = 0.70 \leq F_2(B) \leq \min(1.20 + 1.65 \times 0.12, 0.60) = 0.60$. This is still infeasible, so the second iteration yields $\max(0.70 - 1.65 \times 0.12, 0) = 0.50 \leq F_2(B) \leq \min(0.60 + 1.65 \times 0.12, 0.60) = 0.60$, at which point the constraint becomes feasible and relaxation stops. The upper bound is not relaxed due to it already being at max.

3.2.3 *Bound relaxation after slider interaction.* Since slider adjustments may re-introduce infeasibility, we re-relax the bounds afterwards using Algorithm 1. Once feasible bounds are obtained, we snap the slider to the position whose bounds are closest to the feasible bounds and subsume them, ensuring transparency to the user. Unlike the initial relaxation, we have an advantage here since we only need to relax bounds for a single constraint, with all other constraint bounds fixed. Therefore, if previously feasible bounds (lb_j^*, ub_j^*) are known, we use them as an additional reference during relaxation. Specifically, we reduce lb_j until $\max(0, lb_j^*)$ (line 10 of Algorithm 1) and increase ub_j until $\min(ub_j^*, F_j(T^q))$ (line 11). Additionally, if the user requests termination of relaxation before the process finishes, we move the slider to the closest feasible bounds.

3.3 Devising a quality function for optimality

As discussed in §2.1.2, the quality function defines the notion of optimality used to break ties among multiple valid bundles. This becomes particularly important as bounds relaxation often yields multiple valid bundles. The quality function is application-specific and is typically specified by a domain expert prior to user interaction with an instance of Ex2BUNDLE, as shown in step (0) in Fig. 5. Given (i) a tuple-level scoring function $\sigma : t \mapsto \mathbb{R}$, which assigns a numeric score to a tuple t , and (ii) an application-specific aggregate function $\mathcal{A}$, the quality of a candidate bundle B , $quality(B) = \mathcal{A}_{t \in B} \sigma(t)$. Below, we present example quality functions for the three applications discussed in §1.

Supplier selection for business. Cost and reliability are common quality functions when selecting a bundle of suppliers for business expansion into a new country. For cost, a natural aggregate is SUM, with the objective of minimizing total cost, i.e., $\sum_{t \in B} \text{cost}(t)$. For reliability, suitable aggregates include AVG or MIN, depending on the preference: maximizing average reliability, or maximizing the minimum reliability among the selected suppliers.

Focused snippet extraction. Informativeness and conciseness are common quality functions when extracting snippets from a text document. Sentence-level informativeness can be modeled using term frequency within a sentence t , weighted by inverse document frequency in the corpus (TF-IDF), which assigns higher scores to sentences containing rarer words w.r.t the corpus. Informativeness can also be approximated using structural signals such as the presence of numbers, citations, or hyperlinks, which often indicate factual or reference-rich content. Conciseness can be modeled using the word or character count of a sentence, with the objective of minimizing the total number of words or characters in a snippet. In both cases,

the most appropriate aggregate is SUM. Alternatively, conciseness can be modeled by minimizing the COUNT of sentences in a snippet. *Playlist recommendation.* Popularity and diversity are common quality functions in playlist recommendation. Popularity can be modeled using the number of weekly plays or chart rankings of a song, with the objective of maximizing overall popularity using aggregate AVG, i.e., $\text{AVG}_{t \in B} \text{popularity}(t)$. Diversity can be modeled using pairwise dissimilarity between songs based on features such as genre, mood, or artist. This can be captured via a distance function over song features, with the objective of maximizing diversity within the playlist. A suitable aggregate is MIN over all pairwise distances, i.e., $\text{MIN}_{t_i, t_j \in B, i \neq j} \text{distance}(t_i, t_j)$, and maximizing this ensures that even the most similar pair of songs in the playlist remains sufficiently dissimilar. While in this paper we assume a linear, tuple-wise scoring function, Ex2BUNDLE can support quadratic objectives via CPLEX’s quadratic programming or MILP linearization.

3.4 Package query execution

For the package query execution engine, we use PaQL [12], a query engine for declarative package queries. PaQL translates queries into Integer Linear Programs (ILPs), which are then solved using off-the-shelf solvers such as IBM CPLEX [41]. For large target data with many tuples, solving the ILP becomes computationally expensive due to the combinatorial nature of the search space. To address this, Ex2BUNDLE first PREFILTERS T^q to at most $10x$ tuples most similar to the examples, where x is the average example size; the ILP then enforces the constraints and selects the optimal bundle from that pool. Since this reduction shrinks as x grows, Ex2BUNDLE also applies PaQL’s SKETCHREFINE [12], which solves the ILP over cluster representatives with a guaranteed $(1 + \epsilon)^6$ approximation of the optimal bundle. Empirical results are in §4.5.

Ex2BUNDLE formulates a PaQL query that encodes the constraint bounds Θ (§3.2) as constraints over packages (analogous to snippets in FTSE or bundles in general) and defines the optimization objective using the scoring function σ together with the aggregate $\mathcal{A}$ (§3.3), as shown in §2.2. It then executes the package query, and, if feasible, returns the optimal bundle. The PaQL engine also informs the bound synthesizer (§3.2) about any violated constraints when the query is infeasible in line 5 of Algorithm 1.

EXAMPLE 7. For the constraint bounds $\Theta = \{ \langle 1.50, 2.70 \rangle, \langle 0.50, 0.60 \rangle, \langle 0.10, 0.40 \rangle \}$ (derived in Example 5 and relaxed in Example 6) and the target document T_{CA}^q , we want the snippet with the minimum total word count. The function $\text{word_count} : t \mapsto \mathbb{N}^+$ returns the number of words in a sentence t . The synthesized PaQL query is as follows:

```
PQ( $T_{CA}^q$ ,  $\Theta$ ): SELECT PACKAGE(*) AS B FROM  $T_{CA}^q$ 
  SUCH THAT  $F_1(B)$  BETWEEN 1.50 AND 2.70
  AND  $F_2(B)$  BETWEEN 0.50 AND 0.60
  AND  $F_3(B)$  BETWEEN 0.10 AND 0.40
  MINIMIZE SUM $_{t \in B}$  word_count( $t$ )
```

Remark: Example-based inference assumes the provided examples are correct, yet in practice they may be noisy or insufficiently representative. Ex2BUNDLE mitigates this by letting users adjust the intent sliders directly. Alternatively, examples that differ substantially from the others could be flagged for the user to verify, as in prior work on disambiguating user intent in programming by example [58]; we leave this to future work.

4 EXPERIMENTAL RESULTS

We now present experimental results addressing two sets of research questions. RQ1–RQ3 evaluate to what extent Ex2BUNDLE’s PaQL synthesis algorithm meets the three *correctness criteria* for BQbE: constraint satisfaction, alignment with user intent, and consistency between constraints and bundles (detailed in §4.1). RQ4–RQ6 focus on scalability, practical efficiency, and robustness aspects.

RQ1 To what extent do the bundles returned by Ex2BUNDLE satisfy the synthesized constraints? How does it compare to alternatives? (§4.2).
 RQ2 To what extent do the bundles returned by Ex2BUNDLE align with user intent? How does it compare to LLM-based baselines? (§4.3.1).
 RQ3 Does moving farther from the synthesized constraints produce increasingly less aligned bundles with the intended bundle? (§4.4).
 RQ4 How does Ex2BUNDLE’s runtime scale with the size of example bundles and the size and dimensionality of the documents? (§4.5).
 RQ5 How often is constraint relaxation triggered as example size increases, and how effective is it for BQbE? (§4.6.1).
 RQ6 How robust is Ex2BUNDLE to varying degrees of skewness in the target document’s topic distribution? (§4.7).

Implementation. Ex2BUNDLE uses a Python backend—with Gensim for frequency-based topic modeling, SBERT [84] for semantics-aware topic modeling, IBM CPLEX 20.1 via docplex for ILP solving, and a JavaScript/Bootstrap frontend for the slider interface. We ran experiments on a shared compute cluster, utilizing 4 Intel Emerald Rapids CPU cores, 64 GB RAM, and an NVIDIA H200 GPU.

Datasets. We evaluate Ex2BUNDLE in two applications: (1) supplier selection over the TPC-H dataset [87], using a synthesized view of Supplier (100 rows), Partsupp (8,000 rows), and Nation (25 rows) at scale factor 0.01, with five attributes: price, availability, balance, region_europe, and region_america; and (2) focused text snippet extraction (FTSE) over two datasets: (a) SUBSUME [95], containing 275 (user, intent) pairs, each with 8 manually curated summaries; we use 5 summaries as examples and the remaining 3 for test; and (b) CNN/DailyMail [35], containing human-written article highlights spanning four news categories. Since these highlights are abstractive, we use ChatGPT-4o to align each with sentences from its source article, capping the resulting extractive snippet at 10 sentences or 30% of the article’s sentences.

Baselines. We consider the following baselines, where k is the average example size:

- *Greedy* greedily selects k tuples by decreasing objective score.
- *Cluster* uses K-means clustering over the attribute space, and selects the tuple with highest objective score from each cluster.
- *Random* selects k tuples uniformly at random.
- *Top- k* selects the k most similar sentences to the examples by SBERT [84] cosine similarity.
- *SuDocu* [22] is our prior FTSE system, using Latent Dirichlet Allocation [7] topic distributions.
- *BERTSUMEXT* [54] is an extractive summarizer adapted for BQbE by pre-filtering sentences based on similarity to the examples.
- *MEMSUM* [27] is a reinforcement-learning extractive summarizer for long documents, adapted via the same pre-filtering.
- *ChatGPT-4o* [67] is prompted for FTSE.
- *LLM-agent* uses Llama-3.3-70B-Instruct-FP8 [59] for agent-based PaQL synthesis (including bounds), relaxation, and execution.
- *LLM-Ex2BUNDLE* uses the same LLM for PaQL synthesis, but with Ex2BUNDLE’s constraint relaxation and execution.

4.1 Correctness Criteria & Metrics

We define the correctness criteria for BQbE along three axes:

Constraint satisfaction provides a formal correctness guarantee for feasible queries by ensuring that the returned bundle satisfies the synthesized constraints. To measure the degree of constraint satisfaction, we use *Constraint Satisfaction Rate (CSR)*, the fraction of synthesized constraints satisfied by the retrieved bundle.

Intent alignment is a key correctness criterion that ensures the result bundle produced by a BQbE system closely matches the ground truth in terms of user intent. We use application-specific metrics to measure intent alignment: (1) *Objective score*: composite of TPC-H attributes: sum of price, availability, and balance; (2) *ROUGE-1/2/L* [50]: F_1 scores measuring unigram, bigram, and longest-common-subsequence overlap between result and ground truth; and (3) *Semantic Similarity (SS)*: cosine similarity between average SBERT embeddings of result and ground-truth snippets.

Constraint-bundle alignment consistency evaluates whether deviations from the synthesized PaQL Q produce corresponding deviations in the resulting bundle. Specifically, as a query Q' moves farther from Q , its resulting bundle should become increasingly less aligned with the intended bundle. We measure this using Pearson’s correlation coefficient (PCC) between the Hausdorff distance [72] between the feasible regions induced by two queries and the cosine distance between the SBERT embeddings of the resulting bundles.

4.2 Constraint satisfaction effectiveness

We evaluate RQ1 (constraint satisfaction) for the task of supplier selection on TPC-H. We construct 3 example supplier bundles representing diverse procurement strategies (conservative, price-focused, and balanced); Ex2BUNDLE infers constraints from these examples per §3.2, synthesizes a PaQL query, and executes it to retrieve a result bundle. We compare Ex2BUNDLE against 3 baselines: constraint-agnostic *Greedy* (top- k highest-scoring tuples), *Cluster* (best tuples from each of the k clusters), and *Random* (random k tuples).

Table 8 shows that Ex2BUNDLE satisfies all constraints (CSR = 100%) with a competitive average objective score (11.72). Greedy achieves the highest objective score (15.61) but satisfies on average only 33.3% of the inferred constraints, indicating that objective-only optimization fails to enforce constraints. Cluster outperforms greedy (CSR = 66.7%); however, it fails to satisfy many constraints. Random has a higher CSR (70.0%) than Greedy and Cluster, but with a much lower objective score (10.48), highlighting the importance of jointly optimizing both constraint satisfaction and objective.

- 💡 Greedy and Cluster have lower constraint satisfaction rate (CSR) while Random has low objective score.
- 💡 Ex2BUNDLE achieves 100% CSR with a competitive objective score as it accounts for both objective and constraints.

System	CSR (%) ↑	Average Objective Score ↑	Runtime (s) ↓
Random	<u>70.0</u>	10.48	0.0005
Greedy	33.3	15.61	<u>0.0024</u>
Cluster	66.7	<u>14.04</u>	0.4712
Ex2BUNDLE	100.0	11.72	0.0207

Table 8: Ex2BUNDLE achieves full constraint satisfaction with a competitive objective score. Bold = best; underlined = second-best.

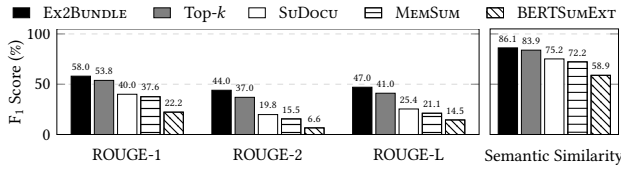


Figure 9: For focused text snippet extraction on the SUBSUME dataset, Ex2BUNDLE outperforms all other baselines across all metrics.

Intent	System	Semantic Similarity ↑
Intent 1: What about this state’s arts and culture attracts you the most?	ChatGPT-4o	0.64
	Ex2BUNDLE	0.88
Intent 2: What are some of the most interesting things about this state?	ChatGPT-4o	0.65
	Ex2BUNDLE	0.76

Table 10: Ex2BUNDLE outperforms ChatGPT-4o on both SUBSUME intents; the advantage narrows on the more general one.

Method	R1	R2	RL	SS	Time (s)	#Relaxation	#Tokens
LLM-agent	0.35	0.17	0.23	0.68	923.60	2.11	104,127
LLM-Ex2BUNDLE	0.39	0.22	0.27	0.74	29.11	2.67	6,536
Ex2BUNDLE	0.56	0.43	0.48	0.85	0.84	1.67	0

Table 11: Ex2BUNDLE outperforms LLM-agent and LLM-Ex2BUNDLE.

4.3 Alignment with user intent

4.3.1 Comparison with non-LLM baselines. For RQ2, we use the SUBSUME dataset and evaluate Ex2BUNDLE against Top-k, SuDocu, BERTSUMEXT, and MEMSUM. We report ROUGE-1/2/L F_1 scores and SBERT-based semantic similarity (SS) between retrieved and ground-truth snippets, together with per-query runtime.

Fig. 9 shows that Ex2BUNDLE consistently outperforms all baselines on ROUGE and SS. SuDocu is a BQbE system and outperforms non-BQbE extractive summarizers MEMSUM and BERTSUMEXT, but its topic modeling is inferior to Ex2BUNDLE and Top-k, both of which use semantics-aware topic modeling.

- 🔥 Ex2BUNDLE outperforms all baselines in terms of ROUGE & SS.
- 🔥 Systems that use semantics-aware topic modeling (Ex2BUNDLE & Top-k) outperform SuDocu and non-BQbE extractive summarization (BERTSUMEXT & MEMSUM).

4.3.2 Comparison with LLM-based baselines as direct snippet extractor. We compare Ex2BUNDLE against LLM-based baselines prompted to directly retrieve snippets from a target document given a few example snippets. We use ChatGPT-4o as a RAG baseline that performs few-shot learning from the examples. We evaluate two intents from SUBSUME—one more specific and one less specific—to assess how each approach handles varying intent specificity (Table 10). Ex2BUNDLE outperforms ChatGPT-4o on both intents, but the advantage narrows for the less specific intent. For generic intents over CNN/DailyMail, ChatGPT-4o outperforms Ex2BUNDLE, which highlights that generic intent is too broad for example-driven retrieval, making Ex2BUNDLE inappropriate. However, the benefit of Ex2BUNDLE being LLM-free is that its latency scales with the corpus rather than model size, it applies to private data where LLM pipelines cannot run, and it incurs no per-query inference cost.

4.3.3 Comparison with LLM-based baselines as intermediate PaQL generator. We additionally compare against two LLM-based baselines that generate PaQL queries as the mechanism for snippet retrieval, rather than directly retrieving snippets. LLM-agent is an

agent-based approach that performs the end-to-end pipeline of PaQL construction, bound relaxation, and query execution, while LLM-Ex2BUNDLE is a hybrid approach that delegates relaxation and execution to Ex2BUNDLE after the LLM generates the PaQL query. We evaluate both approaches on three SUBSUME instances, allowing up to three self-critique-based relaxation steps per instance for LLM-agent. Table 11 summarizes the results.

Ex2BUNDLE outperforms both baselines across all metrics: (1) it is several orders of magnitude faster, (2) it retrieves more intent-aligned bundles, (3) it requires fewer relaxations, and (4) incurs no token cost. LLM-Ex2BUNDLE is 32 $\times$ faster than LLM-agent and outperforms it in terms of quality metrics, runtime, and token cost, but requires more #relaxations. These results highlight two advantages of Ex2BUNDLE: superior initial bound synthesis and a data-aware relaxation strategy that outperforms LLMs’ self-correction loop. Notably, LLM-Ex2BUNDLE’s advantage over LLM-agent persists even when both start from LLM-generated bounds, demonstrating the benefit of Ex2BUNDLE’s data-aware relaxation strategy. In contrast, LLM-generated bounds require more relaxation rounds, increasing not only runtime but also token cost.

- 🔥 Ex2BUNDLE achieves the highest intent alignment across traditional and LLM-based baselines, especially for focused intents.
- 🔥 LLM-generated PaQL queries struggle to produce intent-aligned bundles, while LLM-based bound relaxation is less effective than Ex2BUNDLE’s data-aware relaxation.

4.4 Constraint-bundle alignment consistency

To evaluate RQ3, we test whether larger deviations from synthesized constraints produce larger deviations in the resulting bundles. For each query, we shift all topic constraints by the same additive amount s ($lb' = lb + s$, $ub' = ub + s$), using 15 log-spaced values in $[0.01, 20]$. We retain only feasible, unrelaxed original–shifted pairs, yielding approximately 411 pairs per topic (4,115 total); the resulting Hausdorff distances range from 0.01 to 11.3. We compute PCC between the Hausdorff distance of the induced feasible regions and the cosine distance between the corresponding bundle embeddings. Across topics, the distances are strongly correlated with an average PCC of 0.93, indicating that larger constraint deviations consistently produce larger bundle deviations.

- 🔥 Bounds synthesized by Ex2BUNDLE consistently align with the intended bundles, achieving a PCC of 0.93.

4.5 Scalability analysis

We evaluate Ex2BUNDLE’s scalability (RQ4) w.r.t (a) the size of the examples in #sentences; (b) the size of the target document in #sentences; and (c) the dimensionality of the documents in #topics. Ex2BUNDLE achieves practical scalability via the two mechanisms of §3.4, PREFILTER and PaQL’s SKETCHREFINE. For small x , PREFILTER keeps the ILP size tractable and independent of target document size. However, its effectiveness diminishes as x grows or limiting search to only 10 x sentences becomes suboptimal, which is where SKETCHREFINE contributes.

Over 400 instances of the SUBSUME dataset, we vary the example size from 5 to 100 sentences and the target document size from 100K to 1M sentences, using 5 source and 45 target documents per instance. We increase the number of sentences in each document by

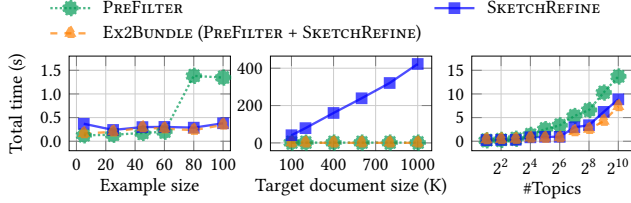


Figure 12: Ex2BUNDLE achieves scalability w.r.t example size (left), target document size (center), and topic dimensionality (right).

paraphrasing the original sentences using an LLM. To synthesize the example bundles, we iteratively sample sentences without replacement to avoid duplicates, randomly selecting sentences across topics until the required bundle size is reached.

Fig. 12 (left) shows that for small documents (≈ 500 sentences), SKETCHREFINE scales better than PREFILTER as #sentences in the examples (x) grows. This is because PREFILTER relies on direct ILP: when the document contains fewer than $10x$ sentences, it retains the full target document. In contrast, SKETCHREFINE always considers the full document but reduces the search space via clustering and invokes ILP over cluster representatives. Ex2BUNDLE applies PREFILTER before SKETCHREFINE, benefiting from both. Fig. 12 (center) shows that for a fixed example size ($x = 100$), PREFILTER’s runtime increases linearly with target document size due to the single pass required to retain the top 1000 sentences, while SKETCHREFINE’s runtime increases more steeply. Ex2BUNDLE benefits from both and inherits the scalability of PREFILTER, with runtime under 0.5 seconds even for documents with 1M sentences. Fig. 12 (right) shows that increasing document dimensionality in #topics (or constraints) causes linear runtime growth for all approaches (note x-axis is in log scale). However, practical workloads rarely exceed 1024 constraints, and even then, Ex2BUNDLE keeps its runtime under 9 seconds.

- Ex2BUNDLE scales nearly independently of document size.
- Ex2BUNDLE’s runtime grows linearly with #topics, and remains under 9 seconds even for practically large #constraints.

4.6 Relaxation behavior and alternatives

4.6.1 Relaxation behavior analysis. We analyze relaxation behavior (RQ5) by varying the example size from 10 to 10K sentences and the target document size from 10 to 1M sentences. About 22% of queries are infeasible when varying example size and 47% when varying target document size. For infeasible queries, Fig. 13 (left) shows that #relaxations grows with example size, as more examples yield a higher lower bounds of the constraints, making them tighter. Conversely, #relaxations decreases with target document size, since larger documents are more likely to contain sentences satisfying the example-derived bounds (Fig. 13 (right)).

4.6.2 Alternative relaxation mechanisms. We contrast our symmetric relaxation (SYM), which relaxes both bounds of each violated constraint equally, against two alternatives: (1) *directional* (DIR) relaxation, which uses an Irreducible Infeasible Subsystem to determine each bound’s relaxation direction, and (2) *FeasOpt* [40], which minimizes the total slack needed for feasibility. We force infeasibility using $pad \in \{30\%, 40\%, 50\%, 60\%\}$ with $(lb, ub) \mapsto ((1 + pad) \times lb, (1 - pad) \times ub)$. Across 80 queries, these levels yield median relaxation rounds of 2/4/7/13 and infeasibility rates

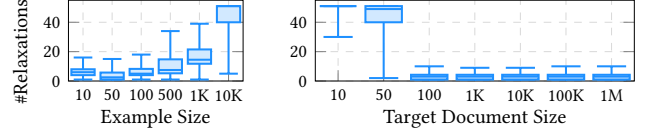


Figure 13: Number of relaxation rounds needed to reach feasibility as (left) example size and (right) target document size grow.

Rounds (n)	1–2 (8)	3–4 (9)	5–6 (21)	7–9 (18)	10–14 (22)	15+ (2)
SYM	1.41	2.16	1.38	1.86	8.59	5.83
DIR	1.48	1.98	1.34	1.88	8.64	5.39
FeasOpt	1.45	1.27	1.52	1.53	1.84	1.33

Table 14: Runtime comparison (in seconds) among Ex2BUNDLE’s relaxation technique SYM, directional relaxation (DIR), & FeasOpt [40].

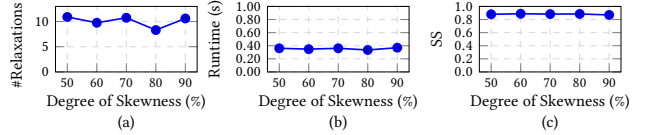


Figure 15: Ex2BUNDLE’s (a) number of relaxations, (b) runtime, and (c) intent alignment as the topic skew in the target document varies.

of 66%/78%/97%/100%. Table 14 (bin size n) shows that FeasOpt outperforms the other methods in runtime mostly at high severity (≥ 10 rounds). However, FeasOpt produces lower-quality bundles: in a separate evaluation over 177 queries, its quality score is lower than SYM’s in 85% of cases and 8% lower on average. SYM matches DIR in 96.6% of cases and is 0.08% higher on average otherwise, since its relaxed feasible region is a superset of DIR’s.

- More examples result in tighter constraints and require more #relaxation; larger documents provide more options to satisfy those bounds, requiring fewer #relaxations.
- The quality of SYM is never lower than DIR’s, while FeasOpt’s quality score is worse than SYM’s in 85% of cases.

4.7 Robustness to data skew

To evaluate robustness to varying degrees of skewness in the target document’s topic distribution (RQ6), we designate one topic as the *focus topic* and vary the fraction of sentences drawn from it (50%, 60%, 70%, 80%, and 90%), with the remaining sentences drawn uniformly from the other topics, over 100 instances of the SUBSUME dataset. Fig. 15 shows the effect of increasing topic skew: (a) shows that the number of relaxations stays stable despite skew; (b) shows that runtime remains largely unaffected by target data skew; and (c) shows that the resulting bundles maintain a consistent level of alignment with the ground truth across all skew levels.

- Ex2BUNDLE’s runtime, intent alignment, and #relaxations remain robust to skew in the target document’s topic distribution.

Additional experimental results are in our technical report [15].

5 INTENT SLIDERS: A USER STUDY

To evaluate Ex2BUNDLE’s slider-based intent-refinement interface for FTSE, we ran a within-subjects study with 20 Amazon Mechanical Turk participants (fluent in English, experienced with document search), who identified their preferred US state for relocation from Wikipedia articles using Ex2BUNDLE versus SBERT, the Top- k baseline (§4.3.1). We exclude non-BQbE baselines (e.g., conversational

LLMs), which underperform on FTSE (Table 10). We set the slider parameter $\alpha = 0.1$ and, following a mixed-methods approach [26], randomized each participant’s starting tool (anonymized) to mitigate transfer effects. Each participant was assigned 25 randomly selected states per tool and asked to provide snippets for at least 2, with no upper limit. With Ex2BUNDLE, participants could select sentences, search by keyword, and refine intent via sliders; with SBERT, they could select sentences and search by keyword too, but had to adjust intent by manually adding or removing example sentences.

We quantitatively analyzed the data by comparing the participants’ ratings of summary quality, ease of use, and preference for both SBERT and Ex2BUNDLE. We also compared the number of interactions with the interface (e.g., selecting examples, adjusting sliders, requesting snippet retrieval, etc.) between the two tools to understand how participants engaged with them. Participants’ open-ended feedback on their experience additionally surfaced signals relevant to the perceived interpretability of the sliders. The majority used the sliders to refine their intents, and their comments suggest an interpretability benefit, for instance: “I used them [sliders] to reduce importance of certain things and increase the importance of others to see how that would affect generated summaries” and “I raised the slider up so that the output summaries would prioritize information that pertained to the economy and climate of each state.”

Participants’ satisfaction. On task completion, participants rated each tool on usefulness, ease of use, and daily usage preference. Ex2Bundle received higher ratings than SBERT on all three: (i) 12/20 rated Ex2Bundle more useful (8/20 for SBERT), (ii) 16/20 found Ex2Bundle easier to use (only 4/20 for SBERT), and (iii) 11/20 preferred Ex2Bundle for daily use (9/20 for SBERT).

Participants’ interactions. To understand how participants interacted with the tools, we categorized interactions into six types: *Search*: keyword queries, *Selection*: example selection, *Update*: example modification, *Learn*: intent learning, *Retrieval*: backend retrieval, and *Slider*: interactive parameters. Averaging interaction counts per user, SBERT users performed about 3.1 times more *Search* interactions than Ex2BUNDLE users. In contrast, Ex2BUNDLE users had more *Slider* interactions and consequently about 2.1 times more *Retrieval* interactions, as Ex2BUNDLE automatically retrieves bundles upon slider adjustments. Ex2BUNDLE’s interactive sliders encouraged more iterative refinement of results, while SBERT’s lack of interactivity led to more manual searching.

6 RELATED WORK

Programming by example [29–31, 58] is a paradigm where users provide examples to express their intent and has seen significant success in data tasks such as data discovery by example [47, 61, 71] and query by example (QbE) [9, 18, 23, 24, 70, 77, 101]. QPlain [18] incorporates explanations with data provenance, some works prune the search space for efficient query synthesis [9, 77] or generalize the output [70], and recently SQvID [23, 24] semantically synthesizes SQL queries from user-provided examples rather than relying on structural similarity. However, all existing QbE systems are designed to retrieve individual items and do not account for the combinatorial nature of bundle retrieval. Ex2BUNDLE addresses the bundle retrieval problem through the lens of QbE, where users express their intent through example bundles, as opposed to example tuples.

Focused text snippet extraction, also known as query-focused extractive summarization, allows users to specify keywords or queries to guide snippet selection [43, 49, 51], but requires users to explicitly formulate their intents. RAG-based approaches [2, 33, 48, 52] can also support snippet extraction via few-shot prompting [11], but are not designed to jointly satisfy multiple constraints. *Faceted query-by-example* approaches leverage representative documents as an example query, additionally conditioned on explicit facets (constraints) to extract snippets across multiple constraints [19, 64, 89]. However, existing systems either require users to specify constraints via keywords [19] or natural language [89], or rely on opaque vector matching that lacks interpretability [64].

Bundle retrieval involves selecting an optimal set of items subject to a set of constraints [12, 14, 20, 57, 91]. It has also been studied extensively in recommender systems across various domains, such as e-commerce [53, 82, 93], entertainment [13, 68], and travel planning [21, 28, 98]. However, these systems rely on greedy heuristics or approximate generative models to generate bundles and struggle to enforce strict constraints. Package queries increase bundle query efficiency [12, 57], but require users to formulate the query. Diversification-based retrieval [14, 20, 91] retrieves bundles that collectively satisfy user constraints, but targets simpler constraints involving cardinality or diversity. In contrast, Ex2BUNDLE addresses the challenge of automatically synthesizing multi-constraint package queries from user examples.

Constraint relaxation. FeasOpt [40] relaxes infeasible constraints to obtain a feasible solution, potentially at the cost of solution quality. Related work in provenance and explanation addresses similar problems but in different settings. CAPE [60] explains surprising aggregate results by mining counterbalancing patterns from data outside the query’s provenance, while QFix [90] diagnoses data errors by computing minimal repairs to historical queries that could cause the observed discrepancy. However, these approaches do not directly apply to our PaQL constraint relaxation problem.

7 SUMMARY AND FUTURE DIRECTIONS

We introduced Ex2BUNDLE, an example-driven framework for bundle retrieval that enables users to specify intents through example bundles. Ex2BUNDLE synthesizes package queries from examples and employs principled constraint relaxation to ensure feasibility. Our experiments and user study show that Ex2BUNDLE captures user intent and outperforms other baselines in intent alignment.

Several directions extend the framework. Manual example selection is burdensome, especially for complex intents or large data; an adaptive example recommendation that infers preference patterns from prior selections would lower this barrier. Ex2BUNDLE currently optimizes linear objectives over per-tuple scores; native PaQL support for pairwise or quadratic objectives (e.g., bundle diversity for playlists) would broaden expressiveness. Specification of the quality function requires domain expertise; inferring both objective and constraints from example bundles—a problem related to inverse optimization [36]—would remove the domain-expertise requirement. Choosing the best aggregate to use in the constraints is another open problem. LLMs can potentially suggest the best aggregate based on the task context, which will lower the configuration burden currently on the domain expert.

REFERENCES

- [1] 2020. Discover Weekly keeps giving me the same genre that I'm completely sick of. <https://community.spotify.com/t5/Your-Library/Discover-Weekly-keeps-giving-me-the-same-genre-that-I-m-d-p/5065068>.
- [2] 2026. LangChain. <https://github.com/langchain-ai/langchain> Accessed: 2026-01-07.
- [3] 2026. Recommendations: Figuring out how to bring unique joy to each member. <https://research.netflix.com/research-area/recommendations>.
- [4] Raghavendra Addanki, Andrew McGregor, Alexandra Meliou, and Zafeiria Moumoulidou. 2022. Improved Approximation and Scalability for Fair Max-Min Diversification. In *25th International Conference on Database Theory (ICDT 2022) (Leibniz International Proceedings in Informatics (LIPIcs))*, Vol. 220. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 7:1–7:21.
- [5] Jinze Bai, Chang Zhou, Junshuai Song, Xiaoru Qu, Weiting An, Zhao Li, and Jun Gao. 2019. Personalized Bundle List Recommendation. In *The World Wide Web Conference, WWW 2019, San Francisco, CA, USA, May 13-17, 2019*. ACM, 60–71.
- [6] Federico Bianchi, Silvia Terragni, Dirk Hovy, Debora Nozza, and Elisabetta Fersini. 2021. Cross-lingual Contextualized Topic Models with Zero-shot Learning. In *Proceedings of the 16th Conference of the European Chapter of the Association for Computational Linguistics: Main Volume*. Association for Computational Linguistics, Online, 1676–1683.
- [7] David M. Blei, Andrew Y. Ng, and Michael I. Jordan. 2003. Latent Dirichlet Allocation. *J. Mach. Learn. Res.* 3 (2003), 993–1022.
- [8] Jay Bonggolto. 2025. Google is testing a new 'Daily Discover' feed on YouTube Music. <https://tech.yahoo.com/streaming/articles/google-testing-daily-discover-feed-170300058.html>.
- [9] Angela Bonifati, Radu Ciucanu, and Slawek Staworko. 2016. Learning Join Queries from User Examples. *ACM Trans. Database Syst.* 40, 4 (2016), 24:1–24:38.
- [10] Virginia Braun and Victoria Clarke. 2012. *Thematic analysis*. American Psychological Association.
- [11] Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. 2020. Language Models are Few-Shot Learners. In *Advances in Neural Information Processing Systems 33: Annual Conference on Neural Information Processing Systems 2020, NeurIPS 2020, December 6-12, 2020, virtual*.
- [12] Matteo Brucato, Juan Felipe Beltran, Azza Abouzied, and Alexandra Meliou. 2016. Scalable Package Queries in Relational Database Systems. *Proc. VLDB Endow.* 9, 7 (2016), 576–587.
- [13] Jianxin Chang, Chen Gao, Xiangnan He, Depeng Jin, and Yong Li. 2023. Bundle Recommendation and Generation With Graph Neural Networks. *IEEE Trans. Knowl. Data Eng.* 35, 3 (2023), 2326–2340.
- [14] Whanhee Cho and Anna Fariha. 2026. Data-Semantics-Aware Recommendation of Diverse Pivot Tables. *Proc. ACM Manag. Data* 4, 1, Article 23 (April 2026), 28 pages.
- [15] Whanhee Cho, Kuangfei Long, Mahmood Jasim, Matteo Brucato, Alexandra Meliou, Peter J. Haas, and Anna Fariha. 2026. Example-Driven Intent Synthesis for Constrained Data Bundle Retrieval: Focused Text Snippet Extraction and Beyond. (2026). [arXiv:2605.19246 \[cs.DB\]](https://arxiv.org/abs/2605.19246) <https://arxiv.org/abs/2605.19246>
- [16] John Joon Young Chung and Eytan Adar. 2023. PromptPaint: Steering Text-to-Image Generation Through Paint Medium-like Interactions. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology (San Francisco, CA, USA) (UIST '23)*. Association for Computing Machinery, New York, NY, USA, Article 6, 17 pages.
- [17] Matthew JC Crump, John V McDonnell, and Todd M Gureckis. 2013. Evaluating Amazon's Mechanical Turk as a tool for experimental behavioral research. *PloS One* 8, 3 (2013), e57410.
- [18] Daniel Deutch and Amir Gilad. 2016. QPlain: Query by explanation. In *32nd IEEE International Conference on Data Engineering, ICDE 2016, Helsinki, Finland, May 16-20, 2016*. IEEE Computer Society, 1358–1361.
- [19] Heejin Do, Sangwon Ryu, Jonghwi Kim, and Gary Lee. 2025. Multi-Facet Blending for Faceted Query-by-Example Retrieval. In *ACL*. 28577–28590.
- [20] Marina Drosou and Evaggelia Pitoura. 2012. DisC diversity: result diversification based on dissimilarity and coverage. *Proc. VLDB Endow.* 6, 1 (2012), 13–24.
- [21] Shih Hsin Fang, Eric Hsueh-Chan Lu, and Vincent S. Tseng. 2014. Trip Recommendation with Multiple User Constraints by Integrating Point-of-Interests and Travel Packages. In *IEEE 15th International Conference on Mobile Data Management, MDM 2014, Brisbane, Australia, July 14-18, 2014 - Volume 1*. IEEE Computer Society, 33–42.
- [22] Anna Fariha, Matteo Brucato, Peter J. Haas, and Alexandra Meliou. 2020. Su-Docu: Summarizing Documents by Example. *Proc. VLDB Endow.* 13, 12 (2020), 2861–2864.
- [23] Anna Fariha, Lucy Cousins, Narges Mahyar, and Alexandra Meliou. 2026. Example-driven semantic-similarity-aware query intent discovery: Empowering users to cross the SQL barrier through query by example. *Inf. Syst.* 138 (2026), 102687.
- [24] Anna Fariha and Alexandra Meliou. 2019. Example-Driven Query Intent Discovery: Abductive Reasoning using Semantic Similarity. *Proc. VLDB Endow.* 12, 11 (2019), 1262–1275.
- [25] Rohit Gandikota, Joanna Materzynska, Tingrui Zhou, Antonio Torralba, and David Bau. 2024. Concept Sliders: LoRA Adaptors for Precise Control in Diffusion Models. In *Computer Vision - ECCV 2024 - 18th European Conference, Milan, Italy, September 29-October 4, 2024, Proceedings, Part XL (Lecture Notes in Computer Science)*. Springer, 172–188.
- [26] Jennifer C Greene, Valerie J Caracelli, and Wendy F Graham. 1989. Toward a conceptual framework for mixed-method evaluation designs. *Educational evaluation and policy analysis* 11, 3 (1989), 255–274.
- [27] Nianlong Gu, Elliott Ash, and Richard H. R. Hahnloser. 2022. MemSum: Extractive Summarization of Long Documents Using Multi-Step Episodic Markov Decision Processes. (2022), 6507–6522.
- [28] Qi Gu, Jian Cao, and Yancen Liu. 2022. CSBR: A Compositional Semantics-Based Service Bundle Recommendation Approach for Mashup Development. *IEEE Trans. Serv. Comput.* 15, 6 (2022), 3170–3183.
- [29] Sumit Gulwani. 2011. Automating string processing in spreadsheets using input-output examples. In *Proceedings of the 38th ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages, POPL 2011, Austin, TX, USA, January 26-28, 2011*. ACM, 317–330.
- [30] Sumit Gulwani. 2016. Programming by Examples: Applications, Algorithms, and Ambiguity Resolution. In *Automated Reasoning - 8th International Joint Conference, IJCAR 2016, Coimbra, Portugal, June 27 - July 2, 2016, Proceedings (Lecture Notes in Computer Science)*, Vol. 9706. Springer, 9–14.
- [31] Sumit Gulwani and Prateek Jain. 2017. Programming by Examples: PL Meets ML. In *Programming Languages and Systems - 15th Asian Symposium, APLAS 2017, Suzhou, China, November 27-29, 2017, Proceedings (Lecture Notes in Computer Science)*, Vol. 10695. Springer, 3–20.
- [32] Wes Gurnee and Max Tegmark. 2024. Language Models Represent Space and Time. In *The Twelfth International Conference on Learning Representations, ICLR 2024, Vienna, Austria, May 7-11, 2024*. OpenReview.net.
- [33] Kelvin Guu, Kenton Lee, Zora Tung, Panupong Pasupat, and Ming-Wei Chang. 2020. Retrieval Augmented Language Model Pre-Training. In *Proceedings of the 37th International Conference on Machine Learning, ICML 2020, 13-18 July 2020, Virtual Event (Proceedings of Machine Learning Research)*, Vol. 119. PMLR, 3929–3938.
- [34] Karl Moritz Hermann, Tomáš Kociský, Edward Grefenstette, Lasse Espeholt, Will Kay, Mustafa Suleyman, and Phil Blunsom. 2015. Teaching Machines to Read and Comprehend. In *Advances in Neural Information Processing Systems 28: Annual Conference on Neural Information Processing Systems 2015, December 7-12, 2015, Montreal, Quebec, Canada*. 1693–1701.
- [35] Karl Moritz Hermann, Tomáš Kociský, Edward Grefenstette, Lasse Espeholt, Will Kay, Mustafa Suleyman, and Phil Blunsom. 2015. Teaching Machines to Read and Comprehend. In *Advances in Neural Information Processing Systems*, Vol. 28.
- [36] Clemens Heuberger. 2004. Inverse Combinatorial Optimization: A Survey on Problems, Methods, and Results. *Journal of Combinatorial Optimization* 8, 3 (2004), 329–361.
- [37] Po Hu, Donghong Ji, Chong Teng, and Yujing Guo. 2012. Context-Enhanced Personalized Social Summarization. In *Proceedings of COLING 2012. The COLING 2012 Organizing Committee, Mumbai, India, 1223–1238*.
- [38] Luyang Huang, Shuyang Cao, Nikolaus Parulian, Heng Ji, and Lu Wang. 2021. Efficient Attentions for Long Document Summarization. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Association for Computational Linguistics, Online, 1419–1436.
- [39] Andy Hunt, Marcelo M. Wanderley, and Matthew Paradis. 2002. The Importance of Parameter Mapping in Electronic Instrument Design. In *New Interfaces for Musical Expression, NIME-02, Proceedings, Dublin, Ireland, May 24-26, 2002*. Media Lab Europe, 149–154.
- [40] IBM. 2015. *IBM ILOG CPLEX Optimization Studio: CPLEX User's Manual* (version 12.6 ed.).
- [41] IBM ILOG CPLEX Optimization Studio [n.d.]. IBM ILOG CPLEX Optimization Studio. <https://www.ibm.com/docs/en/icos>.
- [42] Stratos Idreos, Olga Papaemmanouil, and Surajit Chaudhuri. 2015. Overview of Data Exploration Techniques. In *SIGMOD*. ACM, 277–281.
- [43] Hal Daumé III and Daniel Marcu. 2006. Bayesian Query-Focused Summarization. In *ACL 2006, 21st International Conference on Computational Linguistics and 44th Annual Meeting of the Association for Computational Linguistics, Proceedings of the Conference, Sydney, Australia, 17-21 July 2006*. The Association for Computer Linguistics.

- [44] Rahul Jain, Amit Goel, Koichiro Niinuma, and Aakar Gupta. 2025. AdaptiveSliders: User-aligned Semantic Slider-based Editing of Text-to-Image Model Output. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems (CHI '25)*. Association for Computing Machinery, New York, NY, USA, Article 541, 27 pages.
- [45] Chris Kedzie, Kathleen McKeown, and Hal Daumé III. 2018. Content Selection in Deep Learning Models of Summarization. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, Brussels, Belgium, 1818–1828.
- [46] Shahedul Huq Khandkar. 2009. Open coding. *University of Calgary* 23, 2009 (2009).
- [47] Akash Khatri, Mahathir Mohammad, and El Kindi Rezig. 2025. Sort it Like You Mean It: Discovering Semantically Interesting Attribute Augmentations to Sort Tables. *Proc. VLDB Endow.* 18, 12 (2025), 5427–5430.
- [48] Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. 2020. Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. In *NeurIPS*.
- [49] Yanran Li and Sujian Li. 2014. Query-focused Multi-Document Summarization: Combining a Topic Model with Graph-based Semi-supervised Learning. In *Proceedings of COLING 2014, the 25th International Conference on Computational Linguistics: Technical Papers*. Dublin City University and Association for Computational Linguistics, Dublin, Ireland, 1197–1207.
- [50] Chin-Yew Lin. 2004. Rouge: A package for automatic evaluation of summaries. In *Text summarization branches out*. 74–81.
- [51] Marina Litvak and Natalia Vanetik. 2017. Query-based summarization using MDL principle. In *Proceedings of the multiling 2017 workshop on summarization and summary evaluation across source types and genres*. 22–31.
- [52] Jerry Liu. 2022. *LlamaIndex*. https://github.com/jerryliu/llama_index
- [53] Xiaohao Liu, Jie Wu, Zhulin Tao, Yunshan Ma, Yinwei Wei, and Tat-Seng Chua. 2025. Fine-tuning Multimodal Large Language Models for Product Bundling. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining, V.1, KDD 2025, Toronto, ON, Canada, August 3-7, 2025*. ACM, 848–858.
- [54] Yang Liu and Mirella Lapata. 2019. Text Summarization with Pretrained Encoders. (2019), 3728–3738.
- [55] Xuan Lu, Sifan Liu, Bochao Yin, Yongqi Li, Xinghao Chen, Hui Su, Yaohui Jin, Wenjun Zeng, and Xiaoyu Shen. 2025. MultiConIR: Towards Multi-Condition Information Retrieval. In *Findings of the Association for Computational Linguistics: EMNLP 2025*. Association for Computational Linguistics, Suzhou, China, 13471–13494.
- [56] Gang Luo. 2006. Efficient Detection of Empty-Result Queries. In *Proceedings of the 32nd International Conference on Very Large Data Bases, Seoul, Korea, September 12-15, 2006*. ACM, 1015–1025.
- [57] Anh L. Mai, Pengyu Wang, Azza Abouzied, Matteo Brucato, Peter J. Haas, and Alexandra Meliou. 2024. Scaling Package Queries to a Billion Tuples via Hierarchical Partitioning and Customized Optimization. *Proc. VLDB Endow.* 17, 5 (2024), 1146–1158.
- [58] Mikael Mayer, Gustavo Soares, Maxim Grechkin, Vu Le, Mark Marron, Oleksandr Polozov, Rishabh Singh, Benjamin G. Zorn, and Sumit Gulwani. 2015. User Interaction Models for Disambiguation in Programming by Example. In *UIST*. ACM, 291–301.
- [59] Meta. 2024. Llama 3.3 70B Instruct FP8. <https://huggingface.co/meta-llama/Llama-3.3-70B-Instruct>.
- [60] Zhengjie Miao, Qitian Zeng, Boris Glavic, and Sudeepa Roy. 2019. Going Beyond Provenance: Explaining Query Answers with Pattern-based Counterbalances. In *Proceedings of the 2019 International Conference on Management of Data, SIGMOD Conference 2019, Amsterdam, The Netherlands, June 30 - July 5, 2019*. ACM, 485–502.
- [61] Mir Mahathir Mohammad and El Kindi Rezig. 2026. Qualitative Join Discovery in Data Lakes using Examples. *Proc. ACM Manag. Data* 4, 1, Article 68 (April 2026), 28 pages. <https://doi.org/10.1145/3786682>
- [62] Sina Mohseni, Niloofar Zarei, and Eric D. Ragan. 2021. A Multidisciplinary Survey and Framework for Design and Evaluation of Explainable AI Systems. *ACM Trans. Interact. Intell. Syst.* 11, 3-4 (2021), 24:1–24:45.
- [63] Christoph Molnar. 2020. *Interpretable machine learning*. Lulu. com.
- [64] Sheshera Mysore, Arman Cohan, and Tom Hope. 2022. Multi-Vector Models with Textual Guidance for Fine-Grained Scientific Document Similarity. In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Association for Computational Linguistics, Seattle, United States, 4453–4470.
- [65] Arnab Nandi and H. V. Jagadish. 2011. Guided Interaction: Rethinking the Query-Result Paradigm. *Proc. VLDB Endow.* 4, 12 (2011), 1466–1469.
- [66] Ani Nenkova, Sameer Maskey, and Yang Liu. 2011. Automatic Summarization. In *The 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies, Proceedings of the Conference, 19-24 June, 2011, Portland, Oregon, USA - Tutorial Abstracts*. The Association for Computer Linguistics, 3.
- [67] OpenAI. 2023. GPT-4 Technical Report. arXiv:2303.08774 [cs.CL]
- [68] Apurva Pathak, Kshitiz Gupta, and Julian J. McAuley. 2017. Generating and Personalizing Bundle Recommendations on Steam. In *SIGIR*. ACM, 1073–1076.
- [69] Forough Poursabzi-Sangdeh, Daniel G. Goldstein, Jake M. Hofman, Jennifer Wortman Vaughan, and Hanna M. Wallach. 2021. Manipulating and Measuring Model Interpretability. In *CHI*. ACM, 237:1–237:52.
- [70] Fotis Psallidas, Bolin Ding, Kaushik Chakrabarti, and Surajit Chaudhuri. 2015. S4: Top-k Spreadsheet-Style Search for Query Discovery. In *SIGMOD*. ACM, 2001–2016.
- [71] El Kindi Rezig, Anshul Bhandari, Anna Fariha, Benjamin Price, Allan Vanterpool, Vijay Gadepally, and Michael Stonebraker. 2021. DICE: Data Discovery by Example. *Proc. VLDB Endow.* 14, 12 (2021), 2819–2822.
- [72] R. Tyrrell Rockafellar and Roger J.-B. Wets. 1998. *Variational Analysis*. Springer, Berlin.
- [73] Christian A. Scholbeck, Giuseppe Casalicchio, Christoph Molnar, Bernd Bischl, and Christian Heumann. 2024. Marginal effects for non-linear prediction functions. *Data Min. Knowl. Discov.* 38, 5 (2024), 2997–3042.
- [74] Thibault Sellam and Martin L. Kersten. 2013. Meet Charles, big data query advisor. In *Sixth Biennial Conference on Innovative Data Systems Research, CIDR 2013, Asilomar, CA, USA, January 6-9, 2013, Online Proceedings*. www.cidrdb.org.
- [75] Lesia Semenova, Cynthia Rudin, and Ronald Parr. 2022. On the Existence of Simpler Machine Learning Models. In *FAccT '22: 2022 ACM Conference on Fairness, Accountability, and Transparency, Seoul, Republic of Korea, June 21 - 24, 2022*. ACM, 1827–1858.
- [76] Vidya Setlur, Sarah E. Battersby, Melanie Tory, Rich Gossweiler, and Angel X. Chang. 2016. Eviza: A Natural Language Interface for Visual Analysis. In *Proceedings of the 29th Annual Symposium on User Interface Software and Technology, UIST 2016, Tokyo, Japan, October 16-19, 2016*. ACM, 365–377.
- [77] Yanyan Shen, Kaushik Chakrabarti, Surajit Chaudhuri, Bolin Ding, and Lev Novik. 2014. Discovering queries based on example tuples. In *SIGMOD*. ACM, 493–504.
- [78] Yunxiao Shi, Haoning Shang, Xing Zi, Wujiang Xu, Yue Feng, and Min Xu. 2025. Answering Narrative-Driven Recommendation Queries via a Retrieve-Rank Paradigm and the OCG-Agent. In *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics, Suzhou, China, 13181–13202.
- [79] Aviv Slobodkin, Niv Nachum, Shmuel Amar, Ori Shapira, and Ido Dagan. 2023. SummHelper: Collaborative Human-Computer Summarization. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing, EMNLP 2023 - System Demonstrations, Singapore, December 6-10, 2023*. Association for Computational Linguistics, 554–565.
- [80] Spotify Advertising Team. 2020. Five years of discovery and engagement through Discover Weekly. <https://ads.spotify.com/en-US/news-and-insights/five-years-of-discovery-and-engagement-through-discover-weekly/>. Accessed: 2026-03-17.
- [81] Hendrik Strobelt, Albert Webson, Victor Sanh, Benjamin Hoover, Johanna Beyer, Hanspeter Pfister, and Alexander M. Rush. 2023. Interactive and Visual Prompt Engineering for Ad-hoc Task Adaptation with Large Language Models. *IEEE Trans. Vis. Comput. Graph.* 29, 1 (2023), 1146–1156.
- [82] Wenqi Sun, Ruobing Xie, Junjie Zhang, Wayne Xin Zhao, Leyu Lin, and Ji-Rong Wen. 2023. Generative Next-Basket Recommendation. In *Proceedings of the 17th ACM Conference on Recommender Systems, RecSys 2023, Singapore, Singapore, September 18-22, 2023*. ACM, 737–743.
- [83] Hiroaki Takatsu, Takahiro Kashikawa, Koichi Kimura, Ryota Ando, and Yoichi Matsuyama. 2021. Personalized Extractive Summarization Using an Ising Machine Towards Real-time Generation of Efficient and Coherent Dialogue Scenarios. In *Proceedings of the 3rd Workshop on Natural Language Processing for Conversational AI*. Association for Computational Linguistics, Online, 16–29.
- [84] Nandan Thakur, Nils Reimers, Johannes Daxenberger, and Iryna Gurevych. 2021. Augmented SBERT: Data Augmentation Method for Improving Bi-Encoders for Pairwise Sentence Scoring Tasks. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. Association for Computational Linguistics, Online, 296–310.
- [85] Nandan Thakur, Nils Reimers, Andreas Rücklé, Abhishek Srivastava, and Iryna Gurevych. 2021. BEIR: A Heterogeneous Benchmark for Zero-shot Evaluation of Information Retrieval Models. In *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks 1, NeurIPS Datasets and Benchmarks 2021, December 2021, virtual*.
- [86] Louis L. Thurstone. 1974. *A Law of Comparative Judgment* (1st ed.). Routledge, 12 pages.
- [87] Transaction Processing Performance Council (TPC). [n.d.]. *TPC-H: Decision Support Benchmark*. Technical Report. TPC. <http://www.tpc.org/tpch/>
- [88] Joseph Tso, Preston Schmittou, Quan Huynh, and Jibrán Hutchins. 2026. ConstraintBench: Benchmarking LLM Constraint Reasoning on Direct Optimization. arXiv preprint arXiv:2602.22465 (2026).
- [89] Jianyou Wang, Kaicheng Wang, Xiaoyue Wang, Prudhviraj Naidu, Leon Bergen, and Ramamohan Paturi. 2023. DORIS-MAE: scientific document retrieval using multi-level aspect-based queries. , Article 1668 (2023), 16 pages.

- [90] Xiaolan Wang, Alexandra Meliou, and Eugene Wu. 2017. QFix: Diagnosing Errors through Query Histories. In *Proceedings of the 2017 ACM International Conference on Management of Data, SIGMOD Conference 2017, Chicago, IL, USA, May 14-19, 2017*. ACM, 1369–1384.
- [91] Yue Wang, Alexandra Meliou, and Gerome Miklau. 2018. RC-Index: Diversifying Answers to Range Queries. *Proc. VLDB Endow.* 11, 7 (2018), 773–786.
- [92] Zhijie Wang, Yuheng Huang, Da Song, Lei Ma, and Tianyi Zhang. 2024. PromptCharm: Text-to-Image Generation through Multi-modal Prompting and Refinement. In *CHI*. ACM, 185:1–185:21.
- [93] Penghui Wei, Shaoguo Liu, Xuanhua Yang, Liang Wang, and Bo Zheng. 2022. Towards Personalized Bundle Creative Generation with Contrastive Non-Autoregressive Decoding. In *SIGIR '22: The 45th International ACM SIGIR Conference on Research and Development in Information Retrieval, Madrid, Spain, July 11 - 15, 2022*. ACM, 2634–2638.
- [94] Wen Xiao and Giuseppe Carenini. 2019. Extractive Summarization of Long Documents by Combining Global and Local Context. In *EMNLP-IJCNLP*. Association for Computational Linguistics, 3009–3019.
- [95] Nishant Yadav, Matteo Brucato, Anna Fariha, Oscar Youngquist, Julian Killingback, Alexandra Meliou, and Peter Haas. 2021. SubSumE: A Dataset for Subjective Summary Extraction from Wikipedia Documents. In *New Frontiers in Summarization Workshop*. <https://summarization2021.github.io>
- [96] Rui Yan, Jian-Yun Nie, and Xiaoming Li. 2011. Summarize What You Are Interested In: An Optimization Framework for Interactive Personalized Summarization. In *EMNLP. ACL*, 1342–1351.
- [97] Hamed Zamani, Bhaskar Mitra, Everest Chen, Gord Lueck, Fernando Diaz, Paul N. Bennett, Nick Craswell, and Susan T. Dumais. 2020. Analyzing and Learning from User Interactions for Search Clarification. In *SIGIR*. ACM, 1181–1190.
- [98] Jiale Zhang, Mingqian Ma, Xiaofeng Gao, and Guihai Chen. 2024. Encoder-Decoder Based Route Generation Model for Flexible Travel Recommendation. *IEEE Trans. Serv. Comput.* 17, 3 (2024), 905–920.
- [99] Tianyi Zhang, Faisal Ladhak, Esin Durmus, Percy Liang, Kathleen R. McKeown, and Tatsunori B. Hashimoto. 2024. Benchmarking Large Language Models for News Summarization. *Trans. Assoc. Comput. Linguistics* 12 (2024), 39–57.
- [100] Siyan Zhao, Mingyi Hong, Yang Liu, Devamanyu Hazarika, and Kaixiang Lin. 2025. Do LLMs Recognize Your Preferences? Evaluating Personalized Preference Following in LLMs. In *ICLR*.
- [101] Moshé M. Zloof. 1975. Query-by-Example: the Invocation and Definition of Tables and Forms. In *Proceedings of the International Conference on Very Large Data Bases, September 22-24, 1975, Framingham, Massachusetts, USA*. ACM, 1–24.

APPENDIX A: USER STUDY

To evaluate Ex2BUNDLE, we conducted a within-subjects study with 20 participants recruited via Mechanical Turk. Participants were tasked with identifying their preferred state for relocation based on Wikipedia articles, using both a baseline (SBERT) and Ex2BUNDLE. We followed a mixed-methods approach [26] to analyze quantitative and qualitative data from the study. The study protocol was approved by the institutional review board (IRB) of University of Massachusetts Amherst.

A.1 User Study Settings

Participants. We recruited 20 participants (P01-P20) through Amazon Mechanical Turk [17]. Our inclusion criteria stipulated that the participants must reside in North America, where this study was conducted, and be fluent in English. All participants were qualified Amazon Master Workers, demonstrating a high success level across a wide range of tasks. Participants were compensated with \$25.

Procedure. We followed a single-session within-subject study protocol (Fig. A1). Each study session began with collecting participants’ consent. Participants then watched a 5-minute unskippable video tutorial on using the tool before performing their tasks. After the tutorial, participants completed a pre-study questionnaire to gather their information-seeking purpose, practice, tool usage, time spent, strategy, and ratings of the quality of information. Then, participants were randomly assigned to start with either SBERT or Ex2BUNDLE to negate the transfer effect. During the study, participants interacted with an interface, Fig. A2, to perform their tasks. The interface was designed to allow users to input examples, review and modify them, adjust the importance of different topics using sliders only for Ex2BUNDLE, and view the retrieved snippets. We anonymized the tools’ names to obscure which tool each was using.

Participants were asked to perform a decision-making task: to identify which state in the United States (US) they would relocate to and live in if they were offered their dream job there. We provided them with Wikipedia articles on 25 randomly generated US states. Therefore, each participant worked with two different sets of 25 states across the two tools. We instructed them to read and input at least 2 state snippets with no upper limit. They were also free to use the assigned tool however they wanted.

After working with one tool, we asked the participants to perform the same task with the other tool, but this time, Wikipedia articles were provided on 25 other states. This was done to negate potential unwanted bias and transfer effect from one tool to the other. After completing tasks with both tools, we asked them a set of post-study questions to compare SBERT and Ex2BUNDLE regarding their summary quality, ease of use, and extracted snippets. We also asked them if given the choice, which tool they would use to perform their usual information-seeking tasks.

Data Collection and Analysis. We collected both quantitative and qualitative data. Quantitative data include usage logs, interactions, and time spent working with the tool components. We also collected responses to the Likert scale as well as qualitative answers to open-ended questions pre-study, post-task, and post-study questionnaires. We analyzed the quantitative data by comparing the participants’ ratings such as the summary quality, ease of use, and daily usage of the tools. We analyzed the qualitative data using

thematic analysis [10]. Two members of the research team independently coded data from five randomly selected participants using the open coding method [46] using spreadsheets. The coders discussed, resolved disagreements, and consolidated their codes into a codebook containing the representative set of codes. Then the coders used this codebook to individually code the data from all 20 participants, adding new codes as they emerged. Finally, they compiled the codes into themes across five discussion sessions. The themes were discussed and finalized by the research team.⁴

A.2 Pre-study Questionnaire

In the pre-study questionnaire, we asked participants about their information-seeking purpose, practice, tool usage, time spent, strategy, and ratings of the quality of information. We show bar charts of the results for four of the questions in Fig. A3. The majority of participants (11/20) reported frequently searching across multiple documents or websites. In an information-searching session, the majority of participants (10/20) spent 15-30 minutes, while 2/20 spent <15 minutes, and 5/20 spent 30-45 minutes and 2/20 spent 45-60 minutes, or 1/20 60+ minutes. When asked about the unexpected results, 10/20 said rarely, 7/20 said neutral, and 3/20 said often. This suggests that most participants do not often encounter unexpected results during their information-seeking process. Finally, 8/20 rated the quality of the information gathered after their search process as good. Another 6/20 rated average, while 6/20 rated above average. In summary, the pre-study questionnaire results suggest that most participants frequently search across multiple documents or websites, spend a moderate amount of time on information-seeking sessions, and do not often encounter unexpected results. The quality of information gathered is generally rated as good or above average by most participants.

A.3 Post-study Questionnaire

We asked participants to answer the same set of questions after completing their task with each tool: SBERT and Ex2BUNDLE. The questions were designed to evaluate the participants’ experience with each tool, including their interactions, the quality of the extracted snippets, and their preferences for future use.

Post-Task Questionnaire (SBERT). After completing their task with SBERT, participants answered questions about their experience. Out of the 20 participants with valid responses, one participant didn’t submit their responses for the post-task questionnaire for SBERT, so we analyzed the responses from 19 participants. When asked about modifying their example input snippets, 8 modified their example input snippets during the task while 11 did not. Among all participants, 11/20 rated modifying examples as useful or very useful, 5/20 were neutral, and the remaining 3/20 found it somewhat useful or not useful. Regarding ease of modification, 10/20 rated it as easy or very easy, while 5/20 were neutral. For the quality of the extracted snippets, 14/20 participants rated them as good or excellent (9 good, 5 excellent), while 3/20 rated them average and 2/20 rated them poor. When asked whether the extracted snippets helped them make a decision, 14/20 said yes. Only 4/20 reported encountering unexpected information in the

⁴The codes, and themes are available at <https://github.com/kuangfei-long/ex2bundle/tree/main/demo>.

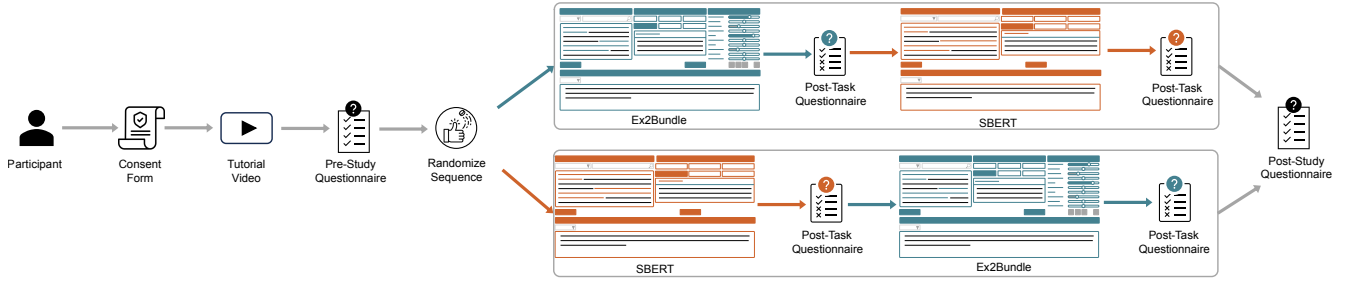


Figure A1: User study protocol. Participants first completed a consent form and watched a tutorial video on how to use the assigned tool. Then, they first performed a pre-study questionnaire, followed by a decision-making task using either SBERT or Ex2BUNDLE. After completing the first task with following post-study questionnaire, they performed the same task with the other tool. Finally, they completed a post-study questionnaire comparing both tools and providing feedback on their experience.

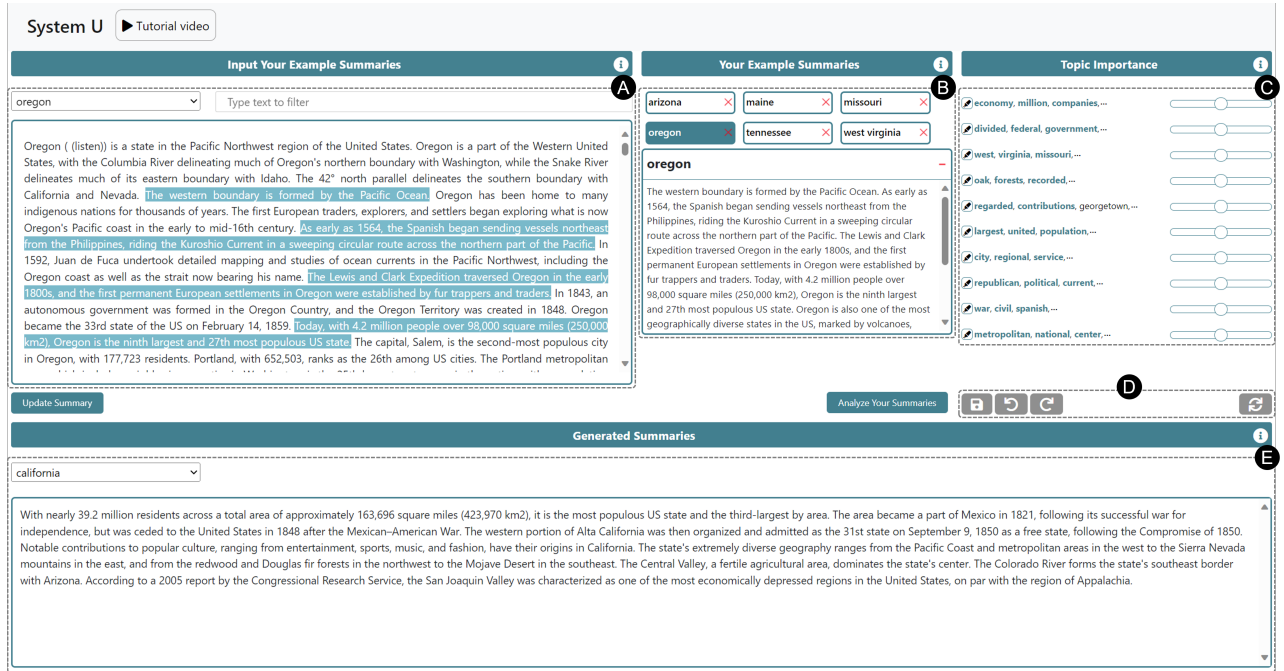


Figure A2: A snapshot of Ex2BUNDLE’s interface. (A) Users can input examples by clicking on the interface, and search for keywords that they are interested in the documents. (B) Users can review their examples and modify them. (C) On the right side, users can adjust the importance of different topics to tune the retrieved snippets using sliders. (D) Users can save their retrieved snippets. (E) The retrieved snippets are displayed under the examples, the figure showing the snippets in California state.

snippets, one participant stating that the snippets showed some information that is not relevant to their input snippets. Regarding future use, 13/20 agreed or strongly agreed that they would use SBERT to search information across multiple documents, while 2/20 disagreed. Participants who liked SBERT highlighted its ease of use and ability to filter relevant content. Common complaints included slow performance, occasional browser crashes, and the absence of importance sliders for tuning results.

Post-Task Questionnaire (Ex2BUNDLE). After completing their task with Ex2BUNDLE, participants answered questions about their experience. We asked them the same questions as in the post-task questionnaire for SBERT, but we also asked them additional questions about their experience with the topic importance sliders, which is a distinctive feature of Ex2BUNDLE. Of the 20 participants with valid responses, 10 modified their example input snippets

while 10 did not. Among all participants, 8/20 found modifying examples useful or very useful, 9/20 were neutral, and 3/20 found it not useful. Regarding ease of modification, 11/20 rated it as easy or very easy, while 7/20 were neutral and 2/20 found it not easy.

A distinctive feature of Ex2BUNDLE is its relative topic importance sliders. The majority of participants (18/20) used the sliders during the task, while 2/20 did not, citing technical issues. Of those who used the sliders, 7/18 observed a change in the extracted snippet after adjusting them. Regarding usefulness of the sliders, 10/20 rated them as useful or very useful, 6/20 were neutral, and 4/20 found them not useful. In terms of ease of use, 15/20 rated the sliders as easy or very easy.

For the quality of the extracted snippets, 12/20 participants rated them as good or excellent (9 good, 3 excellent), while 6/20 average and 2/20 poor or very poor. When asked whether the snippets

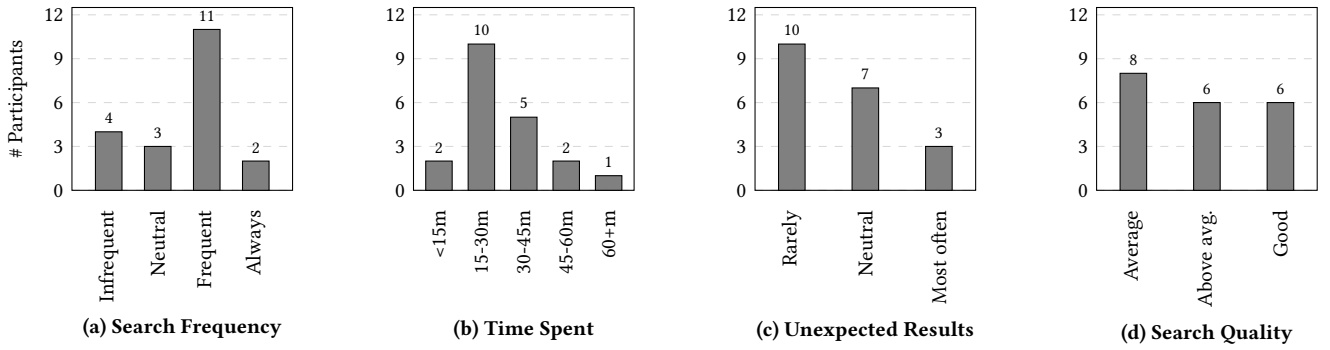


Figure A3: Pre-study questionnaire results (N=20): (a) search frequency from documents or websites, (b) time spent per session (min), (c) frequency of encountering unexpected results, and (d) self-rated search quality.

Question and Response	Count
<i>What was your overall experience of searching for information using summaries?</i>	
Excellent	1
Good	13
Average	3
Very poor	3
<i>While creating example input summaries, did you identify any interesting information that you had not previously thought about?</i>	
Yes	9
No	11

(a)

Which tool did you prefer in each category?	SBERT	Ex2BUNDLE
Which tool generated better summaries?	8	12
Which tool did you find more useful?	9	11
Which tool did you find easier to use?	4	16
Which tool helped you better in making a decision?	8	12
Which tool did you find more flexible?	8	12
Which tool did you find more comfortable to use?	5	15
If you had to choose a tool to search for information from multiple documents/websites, which one would you choose?	9	11

(b)

Figure A4: User study satisfaction results. (a) Overall experience of the user study and discovery of new information. Most participants had a good experience with the system, and more than half of the participants discovered new information. (b) Tool preference responses. The majority of participants preferred Ex2BUNDLE over SBERT across all categories, showing that Ex2BUNDLE was generally favored for extracting text snippets for focused intent tasks.

helped them make a decision, 14/20 said yes. For participants who said no, they mentioned that the snippets were not fully covering the information they wanted. Only 3/20 reported finding unexpected information in the snippets, mentioning that some information was not what they expected based on their input snippets. Regarding future use, 11/20 agreed or strongly agreed that they would use Ex2BUNDLE to search information across multiple documents, while 2/20 strongly disagreed, 3/20 disagreed, and 4/20 were neutral. Participants who liked Ex2BUNDLE highlighted the topic importance sliders, easy navigation, and the ability to customize snippet content. Common complaints included slow performance, freezing, and occasional confusion because of the tool’s technical issues.

A.4 Quantitative Results after Using Both Tools

We first analyzed the quantitative data by comparing the participants’ ratings such as the summary quality, ease of use, and daily

usage preference for both SBERT and Ex2BUNDLE. We also compared the number of interactions of each type between the two tools to understand how participants engaged with them and how much time they spent working with the tools to support their decision-making process.

Participants’ satisfaction. After using both tools, we asked participants to rate the summary quality and ease of use of each tool. In Table A4, overall participants evaluated the experience with retrieving snippets with tools as good (13/20; Table A4 (a)). When comparing SBERT and Ex2BUNDLE, many participants preferred Ex2BUNDLE over SBERT for all dimensions (Table A4 (b)). For example, 12/20 participants preferred Ex2BUNDLE over SBERT for generating better snippets, 11/20 preferred Ex2BUNDLE for usefulness, and 16/20 preferred Ex2BUNDLE for ease of use. When asked which tool they would choose for their usual information-seeking tasks, 11/20 participants preferred Ex2BUNDLE, while 9/20 SBERT.

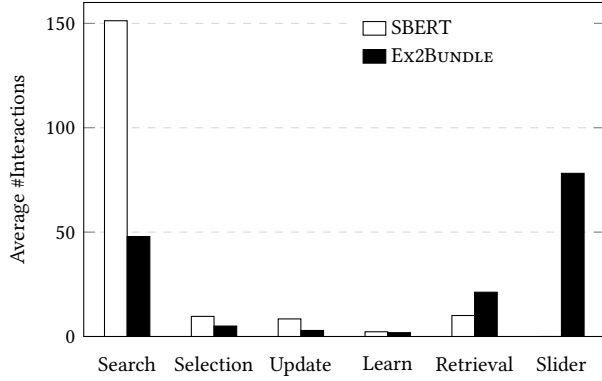


Figure A5: Comparison between SBERT and Ex2BUNDLE in terms of interaction type counts by averaging the number of interactions per user. Here, *Search*: keyword queries, *Selection*: example selection, *Update*: example modification, *Learn*: intent learning, *Retrieval*: bundle retrieval, *Slider*: slider adjustments.

Participants’ interactions. To understand how participants interacted with the tools, we categorized interactions into six types: *Search*: keyword queries, *Selection*: example selection, *Update*: example modification, *Learn*: intent learning, *Retrieval*: backend retrieval, and *Slider*: interactive parameters. *Search* includes keyword queries that participants used to find relevant information in the Wikipedia articles such as “climate”. *Selection* includes interactions where participants selected specific sentences or sections from the Wikipedia articles to include in the input snippets. *Update* includes interactions where participants modified their selected examples, such as adding or removing sentences from the input snippets. *Learn* includes Ex2BUNDLE’s learning process, where it learns the user’s intent based on their example selections. *Retrieval* includes interactions where participants retrieved snippets with PaQL after providing examples or adjusting sliders. *Slider* includes interactions where participants adjusted the importance sliders to tune the retrieval results. We averaged the count of each interaction type per user and compared the counts between SBERT and Ex2BUNDLE.

Fig. A5 shows that SBERT users have more manual interactions: they performed about 3.1 times more *Search* interactions than Ex2BUNDLE users, and selected and updated more examples than Ex2BUNDLE users by about 2.2 and 2.8 times, respectively. However, since Ex2BUNDLE users can adjust the importance sliders unlike SBERT users, they have more *Slider* interactions than SBERT users. Naturally, Ex2BUNDLE users also have about 2.1 times more *Retrieval* interactions than SBERT users, since Ex2BUNDLE retrieves bundles based on slider adjustments, while SBERT does not have such capabilities.

Participants’ querying time. We measure the cumulative backend processing time per participant session. This metric reflects backend request processing time rather than end-to-end participant interaction time. We report this metric to demonstrate the level of engagement and interaction participants had with the system. As shown in Fig. A6, Ex2BUNDLE has a higher average backend processing time than SBERT, with a median of 42.05 seconds for Ex2BUNDLE compared to 17.61 seconds for SBERT. This suggests that participants interacted more frequently and deeply with

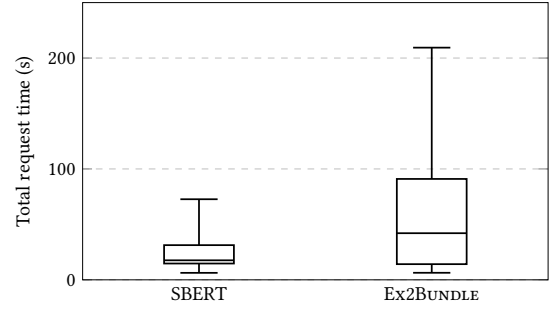


Figure A6: Average request time (s) for SBERT and Ex2BUNDLE. Ex2BUNDLE has a higher average request time than SBERT, which is expected due to the additional processing steps involved in Ex2BUNDLE such as tuning sliders.

Ex2BUNDLE, likely due to its interactive features such as the importance sliders and the ability to retrieve snippets based on user input. In contrast, SBERT’s lower processing time may indicate less frequent interactions, as it does not offer the same level of interactivity and customization as Ex2BUNDLE. Furthermore, on average participants retrieved bundles 15.8 times with Ex2BUNDLE, while only 7.8 times with SBERT, which also supports the notion of higher engagement with Ex2BUNDLE.

A.5 Qualitative Results after Using Both Tools

We analyzed the qualitative data by asking participants to provide feedback on their experience with both tools. We asked them to share their preferences, grievances, and suggestions for improvement. We also asked them to provide feedback on the extracted snippets and how they supported their decision-making process.

Participants preferred Ex2BUNDLE over SBERT. Compared to SBERT, our participants preferred Ex2BUNDLE. When comparing SBERT and Ex2BUNDLE, many participants (at least 11/20 for each category; Table A4 (b)) preferred Ex2BUNDLE over SBERT due to its greater customization. P07 mentioned, “*I absolutely loved the sliders. Being able to rank what is most important to me instead of the system trying to guess is great.*” While performing their tasks in SBERT, several participants (5/20) mentioned that it lacked customizability. P16 said, “*There was no way to adjust which factors were more or less important to me. It was tedious to scroll through long Wikipedia articles looking for relevant info to include, and sometimes articles did not address specific things that I wanted to include. I sometimes filtered/searched for certain words or terms, but there were no results, or none relevant.*” However, some participants (4/20) preferred the simplicity that SBERT provides. P20 mentioned, “*It [SBERT] didn’t require much input to generate summaries that captured what I was most interested in, which made it easier to review the information and make a decision.*”

Topic importance sliders helped to tune snippets. While working on their tasks using Ex2BUNDLE, the participants had access to a set of topic importance sliders to adjust extracted information. The majority of participants (18/20) used the topic importance sliders to fine-tune their snippets and ensure they reflected the information specific to participants’ needs. P08 explained how they used the sliders to focus on the economy and climate for each state, “*I raised*

the slider up so that the output summaries would prioritize information that pertained to the economy and climate of each state since those were the most important factors for me when it came to choosing a state to work in.” In contrast, P09 used the sliders to reduce the importance of certain topics, “Topics that I wasn’t interested in, I just moved to lower importance. [...] That seems to have put more emphasis on certain topics I cared about [in the summary].” The participants (2/20) who did not use the topic importance sliders stated that their choice was predetermined and that changes to the snippets would not have swayed their decision.

Extracted snippet supported informed decision-making. Participants across both conditions (17/20 for SBERT and 20/20 for Ex2BUNDLE) reported that the extracted snippets helped them make informed decisions by enabling them to identify relevant information, gain insight into interesting information, and validate their own preferences. After using Ex2BUNDLE, P04 mentioned, “The summaries show me the details and aspects that I’m interested in while omitting many irrelevant details that I have no interest in knowing. I was able to quickly see which states I find more appealing through reading only the details that I would find relevant.” Similarly, after using SBERT, P07 mentioned, “It was easy to see just the things that were most important to me without the entire Wikipedia page in front of me. Having such detail is nice, but it can be overwhelming. I was able to read the summaries and say, Oh, that doesn’t sound good. I’m going to avoid that, or Awesome! I am going to continue learning more about this state.” Participants were also happy with the quality of the summary and how it organized their preferred data quickly and concisely. P07 pointed out how such efficiencies helped with their tasks, “Being able to switch between the states quickly meant I could see almost instantly how those states compared to each other.”

Participants’ feedback for improvement. The participants suggested several improvements. One of the major grievances of participants was the presence of irrelevant information in the extracted snippet. However, P08 found use in such irrelevant snippet while making decisions: “I read information that showed that North Carolina had a high poverty rate, which I did not remember inputting into the input summaries. However, this data helped me to exclude North Carolina from my list of states to live in.” In addition, some noted the documents’ text-heavy nature. They suggested that, while the snippets were useful for identifying relevant information, creating the input snippets was sometimes tedious and often led to confusion, as going through the Wikipedia article to find candidate sentences for the input was time-consuming. Participants also recommended following the document organization and adding an option to filter sections while reading the articles. Additionally, after using Ex2BUNDLE, 3/20 participants suggested that working with the tool has a steep learning curve. P16 mentioned, “It was a bit confusing at first, even after the tutorial, to know exactly how to use it [Ex2BUNDLE] and what to do for the best results. I did not like having to scroll through really long Wikipedia articles for each state in order to find which parts to highlight and use for the summary.”

APPENDIX B: CONSTRAINT-BUNDLE ALIGNMENT CONSISTENCY

To evaluate whether Ex2BUNDLE’s synthesized bounds produce consistent output, we test whether shifting the synthesized constraint

bounds consistently shifts the returned output bundle. For each query, we shift every topic dimension’s constraint bounds by the same additive amount s (i.e., $lb' = lb + s$ and $ub' = ub + s$), drawn from 15 log-spaced steps in $[0.01, 20]$, avoiding linear spacing since it tends to produce infeasible constraints especially at the extremes. We restrict this experiment to queries whose original and shifted constraints are both feasible and unrelaxed, since a relaxed or infeasible query’s ILP solver falls back to a similar feasible region regardless of the shift and would therefore not test consistency. The Hausdorff distance between the original and shifted bounds is observed to be in the range $[0.01, 11.3]$; shifts producing a larger distance were infeasible. This restriction yields about 411 clean pairs of original and shifted constraints per topic, 4,115 pairs in total across topics. We compute the Pearson correlation between Hausdorff distances of original and shifted constraints and cosine distances of the corresponding output bundles. Across topics, the two distances are strongly correlated (Fig. A7), with an average PCC of 0.93, indicating that different constraint shifts consistently produce different changes in the returned bundle.

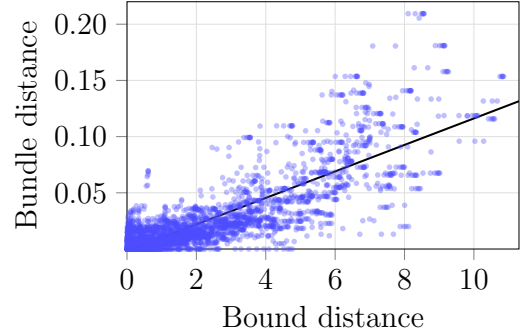


Figure A7: Consistency of constraint shifts and bundle deviations.

APPENDIX C: QUALITY FUNCTION VARIANTS

Fig. A8 demonstrates how different quality function variants affect the performance of Ex2BUNDLE in retrieving sentences on the SUB-SUME. **C-FREQ** is the original implementation introduced in Section 3, which incorporates the term-frequency quality function as the objective function. This variant serves as the base Ex2BUNDLE. **BS** integrates the SBERT cosine-similarity quality function into the Ex2BUNDLE ILP formulation, which maximizes the similarity between selected sentences and the user examples. **P-FREQ** incorporates the TF-IDF and LPR quality functions that capture the importance of sentences based on their frequency and position. **TS** adds a topic-similarity quality function based on cosine similarity, thereby maximizing topic similarity between selected sentences and the examples.

C-FREQ achieves the best performance across all metrics, demonstrating its ability to identify sentences that are semantically aligned with the user examples. Although Ex2BUNDLE relies on this simplified objective function, it outperforms the more complex formulations in terms of both ROUGE scores and semantic similarity. The variant incorporating topic similarity (TS) shows the worst

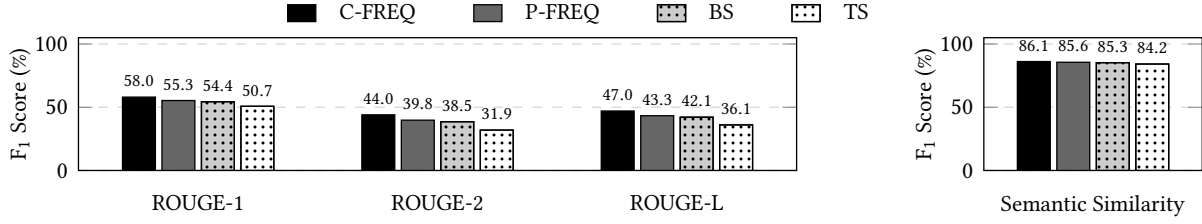


Figure A8: Frequency-based quality functions (C-FREQ & P-FREQ) outperform similarity-based functions (BS & TS) on SUBSUME.

performance across all metrics. This indicates that explicitly maximizing topic similarity is not a reliable signal for selecting sentences that are important and semantically relevant to the examples. Ultimately, while integrating sentence-level semantic similarity or topic similarity seems intuitively promising, they are not effective signals for sentence selection compared to frequency-based objectives. This finding is highly consistent with prior work [45, 85], which observes that frequency-based methods often outperform pure semantic similarity-based methods in retrieval tasks.

- Users can choose different quality functions to optimize for, but the original Ex2BUNDLE formulation with term-frequency quality function (C-FREQ) achieves the best performance across all metrics.
- Frequency-based quality functions outperform semantic similarity-based objectives, suggesting that frequency is a more reliable signal for selecting sentences aligned with user intent in this setting.

collectively, best aligns with user’s information need (focus). For example, when a user wishes to extract sentences from a technical paper that are most relevant to their own research interests. While such intent can, in principle, be expressed in natural language, prior work [22, 95] shows that articulating subjective intents this way is difficult and often leads to user fatigue [97] due to iterative, back-and-forth intent refinement via conversations. Instead, users can provide a few example (document, snippet) pairs to implicitly communicate their focus [37, 79, 83, 96]. For instance, a journalist summarizing reports of all U.S. states with a particular focus—e.g., “basic” economic information, “moderate” coverage of education, and “strong” emphasis on technology—can highlight sentences from a few state’s reports to form example snippets. From these exemplars, Ex2BUNDLE can infer the journalist’s focus, transparently encode it as PaQL constraints, and apply them to new state reports to extract snippets with the same focus. Notably, for reports with substantially different structure or content, Ex2BUNDLE can minimally adjust these inferred constraints to ensure feasible snippet extraction.

APPENDIX D: EXAMPLE USE CASES

We now present three real-world applications where example-driven intent specification lowers the barrier for data bundle retrieval: *supplier selection* for business, *focused snippet extraction* from text documents, and *playlist recommendation* in streaming services.

Supplier selection for business expansion into a new country. Bundle retrieval by example is valuable in business settings where an owner seeks to expand into a new country and must identify a set of new *suppliers* to establish relationships with. The objective is to select a bundle of about 100 suppliers from a large candidate pool over 10,000 suppliers (similar to the TPC-H [87] dataset). Manually constructing an optimal supplier set here is tedious and error-prone, as the selected suppliers must collectively satisfy a range of constraints. Formulating these requirements as an explicit package query is equally challenging: business owners often lack precise knowledge of the exact parameter values for their desired constraints in an unfamiliar market (e.g., acceptable price ranges, inventory availability, or financial stability thresholds). However, they can readily provide example bundles of suppliers drawn from countries where their business already operates successfully. Ex2BUNDLE can leverage these examples to infer the underlying implicit constraints, adjust them to the new supplier database of the target country, and retrieve an optimal set of suppliers that collectively satisfies the constraints.

Focused text snippet extraction. Another use case of Ex2BUNDLE is focused snippet extraction from text documents, also termed as personalized extractive summarization [95]. Here, the goal is to select (extract) a subset of sentences (snippet) from a text document that,

Playlist recommendation for streaming services. In streaming services, users seek recommendations for *bundles* of items that collectively satisfy their preferences. Spotify, YouTube Music, and Netflix exemplify this through playlist/watchlist recommendations (e.g., Discover Weekly [80], Your Daily Discover [8], and Top Picks for You [3]), where the goal is to generate a set of songs/movies aligned with user interest. Existing approaches typically rely on item-level similarity—recommending songs similar to those in a user’s listening history—or on similarity to a single reference playlist. For instance, a user who frequently listens to AC/DC may receive playlists dominated by rock songs, or playlists resembling those curated for users with similar histories. While effective in homogeneous preference settings, such approaches struggle to capture diverse tastes and is notoriously known to cause user frustration, as evident from complaints about lack of variety in Spotify’s Discover Weekly playlist [1]. Consider a user whose taste spans multiple genres—predominantly energetic rock, with a mix of heavy metal, some soft country songs, alongside occasional calm classical music. A playlist based only on rock similarity or a single reference list cannot capture such diversity. Moreover, explicitly specifying such nuanced preferences (e.g., how to balance genre, artist, mood, and tempo) is difficult, even with natural language or LLM-based interfaces. Example-driven bundle recommendation addresses this: users provide a few example playlists that reflect their desired diversity across genres, artists, tempo, and mood, and Ex2BUNDLE infers implicit constraints to generate playlists that collectively satisfy the inferred constraints.

APPENDIX E: DATASET DETAILS

TPC-H. For the package query retrieval task, we use the TPC-H benchmark datasets [87], which consists of a set of business-oriented queries and concurrent data modification. For our experiments, we synthesize a *SupplierFeatures* view derived from three tables using a 0.01 scale factor: *Supplier* with 100 tuples, *Partsupp* with 8,000 rows, and *Nation* with 25 rows. Each supplier is represented by five normalized features: *price competitiveness*, *inventory availability*, *financial stability* (account balance), and two geographic indicators (*America* and *Europe*). Ex2BUNDLE infers sum-based bounds for these features by observing three manually provided example packages representing distinct strategies: Conservative, Price-Focused, and Balanced. The optimization objective is to maximize the aggregate utility of cost, availability, and stability: $Obj = \sum(f_{price} + f_{avail} + f_{bal})$. We compare Ex2BUNDLE against a *Greedy* baseline, which selects the top- k suppliers by raw objective score while ignoring constraints, and a *Random* baseline.

SUBSUME [95]. SUBSUME is a dataset for the evaluation of subjective summary extraction systems. The dataset contains 2,200 (*document*, *intent*, *summary*) triplets over the 50 Wikipedia pages for US states, with ten intents of varying subjectivity, provided by 103 individuals over Mechanical Turk. “Intent” is the underlying question motivating the creation of a summary. In order to evaluate the effectiveness of a *personalized* automatic summarization system, given an intent, we need a few snippets where a subset of the snippets are used as examples provided to the system and the rest are used to evaluate the snippets that the system produces. Accordingly, SUBSUME includes 275 unique (user, intent) pairs, each contributing 8 different, manually curated snippets. During ChatGPT-4o evaluation, five of the snippets are used as example user-inputs for the system to derive a user-intent model in a few-shot manner, and the remaining three are held-back as a test-set to evaluate the system’s ability to capture user-summarization-intent using the five examples. In this work, we explore ten random splits between the example and test sets for each user in the dataset.

CNN/DailyMail. The CNN/DailyMail dataset provides snippets in the form of human-written highlights [34]. However, because these highlights are often abstractive they cannot be directly used for extractive tasks. To adapt this dataset for sentence-level extractive summarization, we employed ChatGPT-4o to align the abstractive snippets with their source articles. The model was tasked with identifying the indices of the article sentences that most closely correlate with the reference highlights. To ensure the resulting extractive snippets remained concise and useful, we imposed two structural constraints: the selection could not exceed 10 sentences or 30% of the original article’s sentence count. These constraints prevent the inclusion of excessive or redundant text, forcing the model to prioritize high-density semantic information.

APPENDIX F: CONSTRAINT EXPRESSIVITY

Here, we provide additional details on how we represent different types of constraints in the integer linear program (ILP) formulation of PaQL. We illustrate three example queries that demonstrate categorical, nonlinear/interaction, and logical operators. Even with linear representations, PaQL can accommodate richer expressions. To validate that these extensions are actually usable by our system,

we ran three case studies on the real-world Spotify Tracks dataset⁵. For each case study we test whether Ex2BUNDLE’s actual bound-synthesis mechanism can formulate a query with the constraint extension, and (2) whether the resulting query can then be solved with Ex2BUNDLE. Concretely, for each case study we froze 5 random example playlists matching a natural-language intent, and solved the resulting package query with an ILP solver:

Q1: Categorical Categorical information can be represented through indicator variables in the integer linear program. To represent this extension, we created a PaQL query about “retrieving a workout playlist that’s mostly dance tracks to keep users energy up, but throw in a couple of hip-hop tracks for variety.” This is a typical example of a categorical constraint as it involves selecting items based on their category (genre=‘dance’). The query below bounds the count of dance tracks with a range, while constraining the count of hip-hop tracks to an exact value, the average number of such tracks across the example playlists:

```
Q1: SELECT PACKAGE(*) FROM SPOTIFY
     SUCH THAT COUNT(*) BETWEEN 10 AND 13
     AND COUNT(genre = 'dance')
           BETWEEN 8 AND 10
     AND COUNT(genre = 'hip-hop')= 3
     MAXIMIZE SUM(popularity);
```

The result shows COUNT(*)=14, COUNT(dance)=11, COUNT(hip-hop)=3, SUM(popularity)=1296. The total and dance-count bounds were infeasible over the real dataset and relaxed.

Q2: Nonlinear/Interaction Feature interactions and nonlinear relationships can be modeled by computing a derived attribute once as a preprocessing step and then bounding it like any other attribute. The intent, “retrieve a quick hits mix that includes tracks that are popular given how short they are, not old epics that only racked up plays because they’re long. About 10 songs, where the average duration is around 3.5 minutes”, includes non-linear attribute hit density. The hit density derived attribute that expresses non-linear relationships is computed as $hit_density = popularity/duration_m$, where $duration_m$ is the track’s duration in minutes and $popularity$ is its popularity score.

```
Q2: SELECT PACKAGE(*) FROM SPOTIFY
     SUCH THAT COUNT(*) BETWEEN 8 AND 12
     AND SUM(popularity / (duration_m))
           BETWEEN 112.0 AND 240.0
     AND AVG(duration_m) = 3.50
     MAXIMIZE SUM(popularity);
```

The result shows COUNT(*)=10, AVG(duration_m)=3.51, SUM(hit density)=239.97, SUM(popularity)=843. The synthesized hit-density bound was infeasible over the real dataset and relaxed upward to reach this optimum.

Q3: Logical Operators Logical conditions can be expressed with AND, OR, and NOT. The intent for Q3 is “retrieve a road-trip playlist that covers a mix of moods: it must include at least one relaxed chill track AND at least one upbeat pop track, but keep every track under 5 minutes. About 10 songs total.” With this intent, we add a logical constraint to ensure the playlist meets the specified criteria.

⁵The Spotify tracks dataset: www.kaggle.com/datasets/maharshipandya/-spotify-tracks-dataset.

Q3:

```
SELECT PACKAGE(*) FROM SPOTIFY
SUCH THAT COUNT(*) BETWEEN 8 AND 12
AND COUNT(genre = 'chill') >= 1
AND COUNT(genre = 'pop') >= 1
AND NOT (COUNT(duration_m > 5) >= 1)
MAXIMIZE SUM(popularity);
```

The result shows COUNT(*)=12, COUNT(chill)=1, COUNT(pop)=11, SUM(popularity)=1132. All synthesized bounds were feasible and no relaxation was needed.

APPENDIX G: BASELINE DETAILS

SBERT [84] is a widely used sentence embedding model that generates semantically meaningful vector representations of sentences. We use SBERT as a simple baseline for example-driven text snippet extraction for selecting sentences that are semantically similar to the example snippets provided by the user based on cosine similarity. We select a snippet using top- k high-scoring sentences in the document, where k is simply the average number of sentences across the examples.

BERTSUMEXT [54] is a widely adopted extractive summarization model. By default, BERTSUMEXT solely considers the input document and extracts important sentences based on general document content, without accommodating specific user inputs. During its training, it learns to identify important sentences by optimizing for word overlap (e.g., ROUGE scores) with gold-standard snippets. To adapt BERTSUMEXT for our example-driven text snippet extraction task, we first pre-filter the target document, retaining only the candidate sentences that are semantically similar to the user-provided example snippets. We then feed this pre-filtered document into BERTSUMEXT. The intuition behind this approach is that by restricting the input space to sentences semantically aligned with the examples, we effectively bias BERTSUMEXT towards extracting a snippet that reflects the user’s specific intent. In our implementation, we use the BERTSUMEXT model pre-trained on the CNN-DailyMail dataset [34].

MEMSUM [27] is a state-of-the-art reinforcement-learning model for long-document extractive summarization. Unlike BERTSUMEXT, MEMSUM is history-aware: it conditions each extraction on the sentences already selected, producing snippets with greater semantic consistency and reduced redundancy. Sentence selection is trained to maximize word overlap with the gold-standard snippet. In this work, we use the model pre-trained on the GovReport dataset [38]. Similarly to BERTSUMEXT, we first pre-filter each target document using the same SBERT-based approach.

ChatGPT-4o [67] is a large language model from OpenAI. We use it as a retrieval-augmented baseline for example-driven text snippet extraction to evaluate how well it can understand user intent and learn extraction patterns from a few examples. We use a file attachment method to provide the model with the target document and example snippets, and we prompt the model to generate a snippet that captures the intent expressed in the examples.

APPENDIX H: PROMPTING DETAILS

SubSumE. For the initial set of ChatGPT-4o few-shot experiments, we conducted inference on each user input file, where each file corresponds to a single intent question. The process involved uploading the relevant Wikipedia text files for the states alongside a

Intent in SUBSUME	ChatGPT-4o Predicted Intent
How is the government structured in this state?	What is the governmental and legislative structure of {state}?
How is the weather of the state?	What is the climate like in {state}?
What drives the economy in this state?	What are the economic drivers and geographical characteristics of different U.S. states?

Table A9: Comparison of original intents in the SUBSUME dataset and intents predicted by ChatGPT-4o using file attachment for context.

structured prompt to instruct the model. For instance, one of the prompts used is as follows:

Few-Shot Prompt for SubSumE

I am providing you with eight text documents. For three of them, Delaware, New Jersey, and Virginia, I will give you example summaries. From these, you need to understand my summarization intent and learn the summarization pattern from the given examples. Then summarize the remaining five documents based on the learned pattern. Make sure to use extractive summarization: select only the original sentences from the given text without generalizing. You can perform detailed summarization by selecting relevant sentences from the text files. Return the final summaries as a combined paragraph.

CNN/DailyMail. For the CNN/DailyMail dataset, we used a similar few-shot prompting approach with ChatGPT-4o as SUBSUME. The prompt was designed to guide the model in learning snippet extraction patterns from a few examples and applying them to new articles. The prompt structure is as follows:

Few-Shot Prompt for CNN/DailyMail

System Instructions:

You are an advanced summarization assistant. I will provide some examples of original articles with its summarization. Please try to learn the pattern of summarization and use it to generate a similar summarization on new given content.

Few-Shot Examples:

{{example_article}}

User Input:

The given new article is {{test_article}}. Please return a list of zero-based sentence indices from the original text that best align with the intent of the summary. Like [1,2].

Output Format:

[index_0, index_1, ...]

APPENDIX I: CHATGPT-4O ON SUBSUME

Here we first report the performance of ChatGPT-4o for predicting intents given example snippets. In Table A9, we show the original intents in the SUBSUME dataset and the intents predicted by ChatGPT-4o via an additional query such as “Given these example summaries, what is the summarization intent?”. The original

Example States	Test State	Semantic Similarity
Task 1: What about this state's arts and culture attracts you the most?		
	Colorado	0.662
Delaware,	New Jersey	0.633
Idaho,	Utah	0.742
Wisconsin	Virginia	0.560
	West Virginia	0.598
	Average	0.639
Task 2: What are some of the most interesting things about this state?		
	Arkansas	0.693
New Hampshire,	California	0.737
Hawaii,	Georgia	0.694
Washington	Kansas	0.533
	Michigan	0.587
	Average	0.649

Table A10: Semantic similarity for ChatGPT-4o across focused intent text snippet extraction tasks, per test state.

		Ex2BUNDLE	ChatGPT-4o (prompt)	ChatGPT-4o (few-shot)
ROUGE-1	Precision	0.2714	0.6699	0.6655
	Recall	0.4769	0.5155	0.5108
	F1	0.3235	0.5504	0.5483
	F2	0.2022	0.3440	0.3427
ROUGE-2	Precision	0.1364	0.5389	0.5304
	Recall	0.2318	0.4106	0.4021
	F1	0.1590	0.4403	0.4332
	F2	0.0994	0.2752	0.2708
ROUGE-L	Precision	0.1977	0.5873	0.5769
	Recall	0.3445	0.4504	0.4405
	F1	0.2335	0.4814	0.4734
	F2	0.1459	0.3009	0.2958
Semantic Similarity		0.6272	0.7965	0.7853

Table A11: Focused text snippet extraction on Ex2BUNDLE, ChatGPT-4o (prompt), and ChatGPT-4o (few-shot) with manually selected ground truth snippet on CNN/DailyMail dataset crime articles.

intents are more specific and focused on particular aspects of the states, while the ChatGPT-4o predicted intents tend to be broader or more general. For example, the original intent of “How is the government structured in this state?” is focused on governmental structure, whereas the ChatGPT-4o predicted “What is the governmental and legislative structure of {state}?”, which is more general and encompasses both governmental and legislative aspects.

Table A10 shows the semantic similarity for ChatGPT-4o across focused text snippet extraction tasks, per test state. For Task 1 (“What about this state’s arts and culture attracts you the most?”), the average semantic similarity across the five test states is 0.639, with Utah having the highest similarity of 0.742 and Virginia having the lowest similarity of 0.560. For Task 2 (“What are some of the most interesting things about this state?”), the average semantic similarity across the five test states is 0.649, with California having the highest similarity of 0.737 and Kansas having the lowest similarity of 0.533. The results give us insights how ChatGPT-4o’s understanding of user intent generally varies across different states,

Category	System	ROUGE-1↑	ROUGE-2↑	ROUGE-L↑	Semantic Similarity↑
Politics	ChatGPT-4o	0.5504	0.4403	0.4814	0.7965
	Ex2BUNDLE	0.3235	0.1590	0.2335	0.6272
Crime	ChatGPT-4o	0.4844	0.3513	0.4019	0.7601
	Ex2BUNDLE	0.3246	0.1499	0.2204	0.5829
Sports	ChatGPT-4o	0.5170	0.4219	0.4548	0.8022
	Ex2BUNDLE	0.3148	0.1743	0.2330	0.6255
Lifestyle	ChatGPT-4o	0.5060	0.3812	0.4294	0.7949
	Ex2BUNDLE	0.3136	0.1304	0.1982	0.6016

Table A12: ChatGPT-4o outperforms Ex2BUNDLE in generic intent alignment across CNN/DailyMail news categories.

which may be influenced by the specific content and characteristics of each state document.

On CNN/DailyMail, we compare against ChatGPT-4o (few-shot) across four categories (Politics, Crime, Sports, & Lifestyle). ChatGPT-4o leads on all ROUGE and SS metrics in every category (Table A12); on Politics, for instance, ChatGPT-4o achieves ROUGE-L 0.4814 and SS 0.7965, vs. 0.2335 and 0.6272 for Ex2BUNDLE (similar gaps elsewhere). This is expected: generic news snippets have broad intent that example-driven retrieval cannot capture [99], marking the regime where Ex2BUNDLE is not the appropriate tool.

APPENDIX J: CNN/DAILYMAIL DETAILS

Here, we provide the detailed results for the generic text snippet extraction and customized focused text snippet extraction tasks on the CNN/DailyMail dataset, as well as the semantic similarity scores for each focused intent text snippet extraction task. For ChatGPT-4o, we report the results for both prompt-only and few-shot prompting approaches. The prompt-only approach uses a single prompt to guide the model’s generation, while the few-shot prompting approach includes three examples for each intent, which are randomly selected from the training set of the CNN/DailyMail dataset.

Generic intent. On CNN/DailyMail, we compare against ChatGPT-4o (few-shot) across four categories (Politics, Crime, Sports, & Lifestyle). ChatGPT-4o leads on all ROUGE and SS metrics in every category (Table A12); on Politics, for instance, ChatGPT-4o achieves ROUGE-L 0.4814 and SS 0.7965, vs. 0.2335 and 0.6272 for Ex2BUNDLE (similar gaps elsewhere). This is expected: generic news snippets have broad intent that example-driven retrieval cannot capture [99], marking the regime where Ex2BUNDLE is not the appropriate tool.

Table A11 shows that ChatGPT-4o (both prompt and few-shot) significantly outperforms Ex2BUNDLE across all ROUGE metrics and semantic similarity, demonstrating its general capability in understanding user intent and generating snippets for generic text snippet extraction tasks. Interestingly, the prompt-only approach performs slightly better than the few-shot approach in terms of ROUGE scores and semantic similarity, which may suggest that providing multiple examples for few-shot prompting does not necessarily lead to better performance in this case.

Table A13 reports the detailed results for the focused text snippet extraction tasks on the CNN/DailyMail dataset contrasting Ex2BUNDLE and ChatGPT-4o (few-shot). Crime articles from CNN/DailyMail contain multi-aspect descriptions (e.g., methods, participants, and locations). However, they are generally written for

broad intent. We compare Ex2BUNDLE on this generic intent against ChatGPT-4o (few-shot). We observe that ChatGPT-4o (few-shot) generally outperforms Ex2BUNDLE across all ROUGE metrics and semantic similarity for the focused intent text snippet extraction tasks, suggesting that ChatGPT-4o is more effective at retrieving snippets aligned with broad user intents than Ex2BUNDLE. However, Ex2BUNDLE still shows competitive performance relative to its own generic-intent baseline in Table A11, suggesting that Ex2BUNDLE remains effective at retrieving snippets for specific intents, even though those intents are still relatively broad.

Tables A14 to A17 report the detailed results for the focused text snippet extraction tasks on the CNN/DailyMail dataset with our manually crafted focused intents for crime, politics, sports, and lifestyle categories. Across all categories, the results show that Ex2BUNDLE’s performance is generally better for the focused intent text snippet extraction tasks compared to the generic snippet extraction task. These findings strengthen the evidence that Ex2BUNDLE is effective in retrieving snippets for specific user intents.

		“Crime methods”		“Suspects & Victims”		“Crime locations”	
		Ex2BUNDLE	ChatGPT-4o (few-shot)	Ex2BUNDLE	ChatGPT-4o (few-shot)	Ex2BUNDLE	ChatGPT-4o (few-shot)
ROUGE-1	Precision	0.5126	0.6831	0.4562	0.6901	0.4016	0.6825
	Recall	0.4489	0.6043	0.4215	0.5855	0.3783	0.4260
	F1	0.4608	0.6206	0.4207	0.5999	0.3686	0.5021
	F2	0.2880	0.3879	0.2629	0.3749	0.2304	0.3138
ROUGE-2	Precision	0.3009	0.5586	0.2493	0.5781	0.2110	0.5010
	Recall	0.2523	0.4928	0.2198	0.4780	0.1862	0.3116
	F1	0.2644	0.5066	0.2235	0.4946	0.1861	0.3684
	F2	0.1653	0.3166	0.1397	0.3091	0.1163	0.2302
ROUGE-L	Precision	0.3526	0.5944	0.3168	0.6079	0.2904	0.5641
	Recall	0.3004	0.5256	0.2850	0.5088	0.2660	0.3498
	F1	0.3124	0.5397	0.2872	0.5242	0.2619	0.4135
	F2	0.1952	0.3373	0.1795	0.3276	0.1636	0.2584
Semantic Similarity		0.7435	0.8307	0.7021	0.8089	0.6117	0.7552

Table A13: Focused text snippet extraction task on Ex2BUNDLE and ChatGPT-4o (few-shot) with LLM selected ground truth snippet on crime detailed description.

		Generic	“Summarize the crime mentioned in this article.”	“Who are the suspects and victims in this event?”	“Summarize the locations where the crime was committed.”
ROUGE-1	Precision	0.2714	0.5126	0.4562	0.4016
	Recall	0.4769	0.4489	0.4215	0.3783
	F1	0.3235	0.4608	0.4207	0.3686
	F2	0.2022	0.288	0.2629	0.2304
ROUGE-2	Precision	0.1364	0.3009	0.2493	0.211
	Recall	0.2318	0.2523	0.2198	0.1862
	F1	0.159	0.2644	0.2235	0.1861
	F2	0.0994	0.1653	0.1397	0.1163
ROUGE-L	Precision	0.1977	0.3526	0.3168	0.2904
	Recall	0.3445	0.3004	0.285	0.266
	F1	0.2335	0.3124	0.2872	0.2619
	F2	0.1459	0.1952	0.1795	0.1636
Semantic Similarity		0.6272	0.7435	0.7021	0.6117

Table A14: Ex2BUNDLE’s performance across various metrics for generic and focused intents for crime category of CNN/DailyMail.

		Generic	“Summarize the politicians and celebrities mentioned.”	“Summarize the contributions or impacts of this political event or statement.”	“What are the key issues or debates highlighted in this political event?”
ROUGE-1	Precision	0.3068	0.4699	0.448	0.5055
	Recall	0.4029	0.397	0.5131	0.4703
	F1	0.3246	0.392	0.4626	0.4684
	F2	0.2029	0.245	0.2891	0.2927
ROUGE-2	Precision	0.1395	0.2807	0.2686	0.2975
	Recall	0.1853	0.2076	0.3019	0.2617
	F1	0.1499	0.2125	0.2751	0.2669
	F2	0.0937	0.1328	0.1719	0.1668
ROUGE-L	Precision	0.2065	0.3422	0.3126	0.3511
	Recall	0.2746	0.2675	0.36	0.3146
	F1	0.2204	0.2699	0.3233	0.3186
	F2	0.1377	0.1686	0.2021	0.1991
Semantic Similarity		0.5829	0.671	0.7454	0.7658

Table A15: Ex2BUNDLE’s performance across various metrics for generic and focused intents for politics category of CNN/DailyMail.

		Generic	“Summarize the sports teams or athletes involved.”	“Summarize the match results or the event details.”	“Where is the stadium or location hosting the sports match event.”
ROUGE-1	Precision	0.3148	0.4544	0.5007	0.3875
	Recall	0.3886	0.5049	0.4372	0.2503
	F1	0.3148	0.4675	0.4488	0.2828
	F2	0.1967	0.2922	0.2805	0.1768
ROUGE-2	Precision	0.1823	0.2903	0.3285	0.2407
	Recall	0.2098	0.3095	0.267	0.1407
	F1	0.1743	0.2932	0.2833	0.1658
	F2	0.1089	0.1833	0.1771	0.1036
ROUGE-L	Precision	0.2405	0.341	0.3872	0.3115
	Recall	0.2852	0.3732	0.3273	0.1863
	F1	0.233	0.3486	0.341	0.2176
	F2	0.1456	0.2179	0.2131	0.136
Semantic Similarity		0.6255	0.7552	0.725	0.6476

Table A16: Ex2BUNDLE’s performance across various metrics for generic and focused intents for sports category of CNN/DailyMail.

		Generic	“What are the main lifestyle trends or topics highlighted in the article?”	“Summarize the key tips or advice offered for enhancing personal well-being and daily life.”	“What significant insights or stories does the article provide about cultural practices or personal experiences?”
ROUGE-1	Precision	0.3056	0.5585	0.4078	0.6109
	Recall	0.3554	0.4031	0.3706	0.4331
	F1	0.3136	0.4533	0.3638	0.4931
	F2	0.196	0.2833	0.2274	0.3082
ROUGE-2	Precision	0.124	0.3291	0.225	0.4043
	Recall	0.1519	0.2235	0.1885	0.2735
	F1	0.1304	0.2596	0.1925	0.3178
	F2	0.0815	0.1622	0.1203	0.1986
ROUGE-L	Precision	0.1919	0.3868	0.2798	0.4454
	Recall	0.2268	0.2726	0.245	0.3032
	F1	0.1982	0.3105	0.2445	0.3514
	F2	0.1239	0.1941	0.1528	0.2196
Semantic Similarity		0.6016	0.7495	0.673	0.7832

Table A17: Ex2BUNDLE’s performance across various metrics for generic and focused intents for lifestyle category of CNN/DailyMail.