
Refactoring Codebases through Library Design

Anonymous Author(s)

Affiliation

Address

email

Abstract

Maintainable and general software allows developers to build robust applications efficiently, yet achieving these qualities often requires refactoring specialized solutions into reusable components. This challenge becomes particularly relevant as code agents become increasingly accurate at solving isolated programming problems. We investigate code agents' capacity to refactor code in ways supporting growth and reusability. We first investigate what makes a good refactoring, finding via a human study that an MDL (Minimum Description Length) objective best aligns with developer preferences for code refactoring quality. We then present both a method and a benchmark for refactoring: LIBRARIAN, a sample-and-rerank method for generating reusable libraries, built on this objective, and MINICODE, a benchmark where code agents must minimize and refactor multiple independent solutions into a joint library. Compared to state-of-the-art code agents, LIBRARIAN achieves strong results on both compression and correctness on MINICODE, obtaining compression rates 1.6-2x better than coding agents while also improving correctness.

1 Introduction

Writing code is mainly a matter of *rewriting* code: debugging, refactoring, optimizing, and other activities within the software engineering lifecycle. But poor rewrites incur technical debt, with such debt costing up to \$2 *trillion* annually [1]. This problem will likely worsen as language models become increasingly responsible for generating code, because they excel at solving isolated programming problems, but their context length demands a myopic view of the codebase. It is therefore valuable to understand not just the ability of language models to solve programming problems, but also their ability to rewrite and refactor code in ways that support growth and reuse.

Effective code refactoring at scale is a design problem. When refactoring codebases, developers must navigate design decisions around concerns such as generality, re-usability, and maintainability. A classic example illustrates this design challenge: Human programmers often create overly-specialized, redundant solutions to similar problems and would benefit from redesigning specialized solutions into a shared library. This consolidation requires careful design decisions about the right level of abstraction — neither too specific nor too general — and appropriate interfaces that balance flexibility with usability.

Here we focus on refactoring multiple code sources into a reusable software library, and pose the following question: To what extent can code agents address this problem, both within human-written codebases, and also in language model-generated code? To answer that question, we develop a new method and a benchmark. This goes beyond past work [2, 3, 4, 5, 6, 7, 8] in *library learning* that synthesized subroutines across small programs in i.e. λ -calculus, instead tackling the more naturalistic problem of redesigning large bodies of code written in contemporary high-level languages, such as Python, producing classes, methods, and helper functions in the style of a human-written library. We develop a method, LIBRARIAN (Figure 1), which samples possible code rewrites then

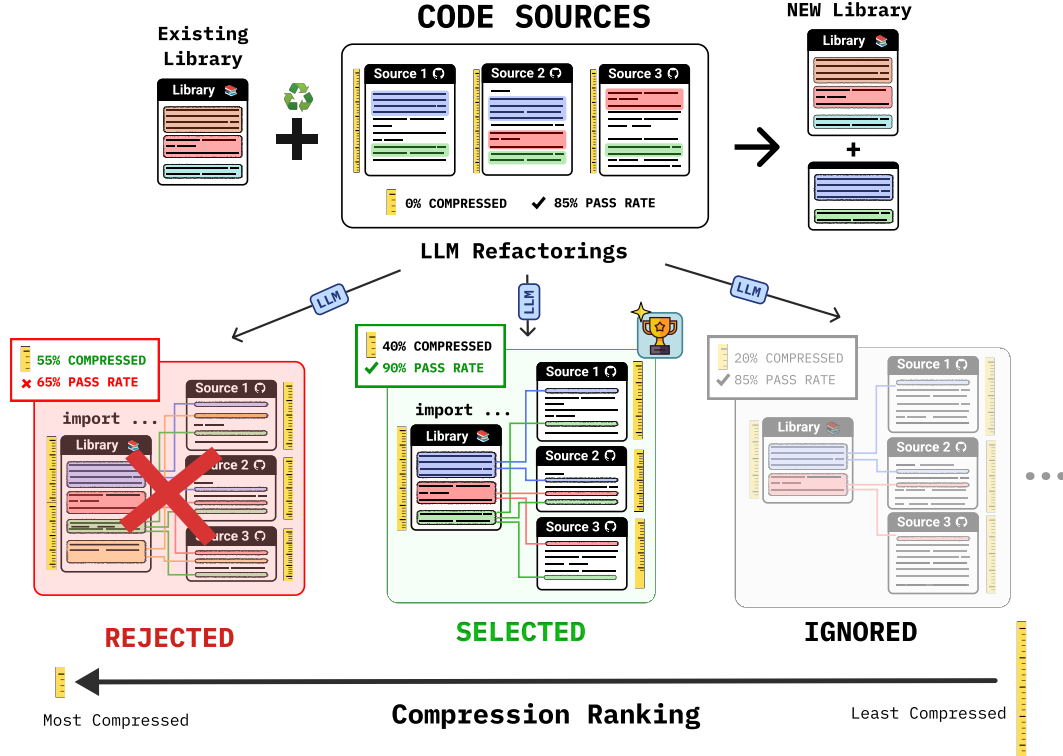


Figure 1: Overview of the refactoring problem and the general structure of its solutions. Given a collection of different code sources, where a source is either program or repository,—and optionally an existing library—we refactor the code sources by designing a new modular and reusable library. Candidate refactorings are evaluated based on program simplicity (compression), and are expected to maintain correctness of the original code sources (pass rate).

39 reranks those samples based on criteria designed to capture what it means to have a good refactoring.
 40 To generate potential rewrites, we develop methods for clustering pieces of code together that share
 41 common structure, so that a succinct prompt can rewrite them jointly into their refactored form.
 42 To find strong criteria for ranking potential rewrites, we study a variety of metrics across machine
 43 learning and software engineering, both on program synthesis benchmarks and via a human study.

44 To evaluate our method and systematically assess the capability of current agents to perform such
 45 design-intensive refactorings, we introduce a new benchmark, MINICODE, which addresses three
 46 key desiderata missing from existing benchmarks. First, open-ended design: unlike SWE-Bench [9],
 47 Commit0 [10], and RefactorBench [11] which primarily focus on functional correctness, MINICODE
 48 presents an unconstrained library design problem. Agents must create a library that can be imported
 49 into multiple repositories, with complete freedom to design the interface and implementation from
 50 scratch—optimizing for software engineering objectives like reusability and maintainability. Second,
 51 verifiability: we ensure objective evaluation by retaining the unit tests from all repositories that will
 52 import the designed library, allowing us to verify that the solutions work correctly across multiple
 53 use cases. Third, large context: agents must understand and synthesize information from multiple
 54 repositories simultaneously to design a unified library that consolidates specialized code sources into
 55 a general interface. Prior benchmarks typically focus on single-repository tasks.

56 Our results show that state-of-the-art code agents, based on Claude Sonnet 3.7, struggle to jointly
 57 preserve correctness and improve reusability across all domains of MINICODE. In the competition
 58 coding domain, our method LIBRARIAN improves refactoring quality by 1.89x while also enhancing
 59 correctness. However, on the repository-level refactoring, even the strongest agents fail to produce
 60 high-quality refactorings, highlighting a substantial gap between current capabilities and the demands

61 of design-oriented code rewriting. Addressing this challenge remains an open and important direction
62 for future research.

63 2 Related work

64 **Library Learning.** Systems which perform library learning research discover shared abstractions
65 across a large number of small programs, which they use to automatically define new subroutines.
66 Systems such as DreamCoder [4], Trove [12], LiLo [7], and REGAL [3] automatically construct such
67 libraries with the goal of making future program synthesis tasks easier to solve, once the learned
68 library is in hand. Our work is closest to REGAL [3], which clusters related code and refactors using
69 language models. However, existing library learning approaches have primarily been demonstrated in
70 small-scale, constrained domains, limiting their applicability to typical software engineering tasks,
71 such as consolidating multiple repositories into cohesive libraries. By framing library learning within
72 the context of realistic, large-scale code repository development, we expand the relevance of library
73 learning to everyday software engineering practice.

74 **Repo-level coding benchmarks.** Recent work has explored the application of language models to
75 repository-level software engineering tasks. Existing benchmarks include SWE-bench [9], which
76 evaluates models on their ability to resolve real-life GitHub issues, and Commit-0 [10], which requires
77 agents to fill in function definitions. Such benchmarks primarily evaluate functional correctness
78 via unit tests, without assessing the quality or maintainability of the resulting codebase. Refactor-
79 Bench [11] takes a step in this direction by benchmarking the ability to follow specific refactoring
80 instructions. Our work differs by requiring models to perform a more open-ended task: Redesigning
81 code to be more modular and compact by discovering and drawing out reused abstractions, while
82 retaining verifiability by re-using downstream unit tests. Additionally, libraries must be created
83 without any scaffolding limitations such as preexisting function definitions affording more design
84 freedom than Commit-0.

85 **Program optimization.** While our goal is to optimize the quality of libraries, other works focus
86 on improving execution speed through correctness-preserving transformations [13, 14, 15]. Both
87 forms of program optimization, compression and speed, are more open-ended than optimizing only
88 for correctness, as there does not exist a ground-truth answer. Prior work on program optimization
89 benchmarks study code at the file level. We propose a benchmark that transforms programs at a larger
90 scale, across multiple code repositories.

91 3 Problem Statement

92 In this section, we propose a refactoring task: Given multiple code sources that contain problem-
93 specific implementations, the goal is to create a cohesive library that captures shared abstractions.
94 This library must reduce the total code size while supporting all original use cases, potentially opening
95 up new use cases as well by mining and formalizing latent shared abstractions. This is accomplished
96 by searching for refactorings that are both correct and simple. Correctness is straightforward to define
97 as the fraction of unit tests passed, but simplicity is more elusive.

98 One potential measure of simplicity is counting the total *number of tokens* in the proposed library
99 and refactored code [6, 16, 5, 17]. However, just minimizing program size has obvious failure modes:
100 code should also be natural, elegant, and extensible, which can be in tension with merely finding the
101 shortest program.¹ Other work in program synthesis [18, 19, 4] defines simplicity as the *minimum*
102 *description length* (MDL), or negative log probability under a reference distribution. In the software
103 engineering community, other metrics such as cyclomatic complexity and maintainability index have
104 been defined for similar purposes: These are more complex metrics that examine the syntax tree, call
105 graph, and other statically-analyzable structures [20]. What metric should we use? The right choice
106 of simplicity metric remains unclear. We return to this question in Section 6, where we empirically
107 compare candidate metrics and human preferences before fixing our choice for the rest of the paper.

¹Perl Golf is a game where participants attempt to write the shortest Perl program accomplishing a given task. The resulting code is famously incomprehensible, even by the standards of Perl.

Without committing to a simplicity metric, let’s use a placeholder metric M ; assume we are given a collection of code sources $\{\rho_n\}_{n=1}^N$, and output both a new library $\mathcal{L}$, as well as rewritten refactorings of the original code sources, $\{\rho'_n\}_{n=1}^N$. We define the pass rate $\tau(\rho_n)$ as the fraction of unit tests program ρ_n passes. In practice we are concerned both with the case where we are refactoring several code sources ($N > 1$) and also the case where there is only a single large code source we are refactoring ($N = 1$).

We optimize the following objective, which rewards refactorings that pass at least as many tests as the original program *and* minimize the chosen metric M :

$$\ell(\mathcal{L}, \{\rho'_n\}) = \begin{cases} M(\mathcal{L}, \{\rho'_n\}) & \forall \rho_n, \tau(\rho_n) \leq \tau(\rho'_n) \\ \infty & \text{otherwise} \end{cases} \quad (1)$$

where $M(\mathcal{L}, \{\rho'_n\})$ may be instantiated in different ways depending on the notion of simplicity under consideration.

One example of M is minimum description length (MDL), given by $M_{MDL}(\mathcal{L}, \{\rho'_n\}) = -\log p_{LM}(\mathcal{L}) + \sum_n -\log p_{LM}(\rho'_n | \mathcal{L})$, where $p_{LM}(\rho'_n | \mathcal{L})$ is concatenating the library and the program into one prompt, but only counting the perplexity of the later program tokens.

4 MINICODE—Library Design and Refactoring Benchmark

MINICODE evaluates a code agent’s capability to identify abstractions across implementations and design reusable libraries. In order to measure these capabilities, our benchmark presents agents with a collection of code sources, then asks agents to refactor the code sources into a unified library alongside refactorings of the original code sources. There are two key desiderata for collections of code sources: The collections must be compressible, in that there exists a latent shared library abstraction, and verifiable, so that we can measure how well refactored code sources preserve functional correctness. See Appendix C for details on our clustering. We source problems from three domains, both real world and synthetic: **Competition coding (Code Contests)**, **Huggingface Transformers** model implementations, and **synthesized repositories** (Table 1).

Agents are expected to interact with MINICODE via the terminal. We structure the benchmark as refactoring a multi-package Python repository, where each code source in a collection is a Python package in a subdirectory. This requires knowledge of basic bash commands for exploring repositories, editing code, and running tests, as well as how to manage complex, multi-package Python libraries.

Table 1: MINICODE Statistics

Domain	Sources	Collections	Avg LoC	Avg Tests	Gen by
Code Contests	300	30	87	10	Humans
Transformers	8	2	538	181	Humans
Synthetic Repositories	20	10	6,433	101	Claude-Sonnet 3.7

CodeContests. Competition problems are crafted with specific variations of algorithmic approaches in mind, resulting in both shared latent concepts and required test cases. As a result, competition coding is naturally both compressible and verifiable.

Each collection consists of multiple code sources, each containing a solution to a competition programming prompt, with associated tests for verification. We take solutions, prompts, and tests from CODECONTESTS [21], a dataset consisting of competition coding problems. Each code source in the collection is structured as a subdirectory consisting of the task description in `PROBLEM.md`, the initial solution in `main.py`, and a script to run tests in `run.sh`. Agents are instructed to create a `library.py` file, which is imported into each code source. Since CODECONTESTS has no external dependencies on Python packages, this can be done without explicit structuring as a Python package.

Huggingface Transformers Library. In this domain, we test whether agents can refactor across implementations of large language and vision-language models (`modelling_<name>.py` files) from the Huggingface transformers repository (e.g., Qwen2, LLaMA, DeepSeek-V3). Unlike

148 competition coding, these sources are production-scale and Huggingface requires that all changes
 149 pass an extensive suite of integration tests before being accepted into the main branch. A refactoring
 150 is only deemed correct if it passes the unmodified Transformers test suite, making this a high-stakes
 151 setting that requires correctness and compatibility.

152 **Synthetic Repositories.** We synthesize repositories by first generating project ideas and then
 153 specialized variations, allowing us to control both complexity and overlap between code sources in a
 154 collection. Each code source in the repositories domain consists of a task description, source code,
 155 and test cases for functionality and correctness. MINICODE includes repositories, approximately
 156 6.5k lines of source code each, which represent realistic settings where different people with different
 157 needs use language models to help them write software for their particular use cases. The refactoring
 158 agent is tasked with extracting re-usable functions from across code repositories, and re-writing the
 159 original code source repositories to use them.

160 Concretely, we prompt coding agents to generate a list of ideas. For each idea, agents generate a
 161 collection consisting of two variations of that idea, targeting an imagined persona. Each variation is
 162 then instantiated in code as a Python repository with at least 100 tests. The prompts for each step
 163 are shared in Appendix F. Each collection is structured as a multi-package Python library: every
 164 synthetic repository becomes a subpackage with its source and tests. Agents are instructed to write a
 165 shared library in a shared subpackage named `common`. This common shared library must be imported
 166 and used to refactor each of the original subpackages.

167 5 LIBRARIAN: Refactoring Code to Create Libraries

168 This section details our method to compress collections of code sources into libraries, while migrating
 169 the code sources to use these shared building blocks. Figure 1 illustrates our method, LIBRARIAN.

170 LIBRARIAN follows a simple sample-and-rerank framework to maximize our refactoring objective
 171 described in Section 3. It maintains and grows a library of useful functions as part of this objective.
 172 Concretely, our framework follows:

$$\mathcal{L}^*, \{\rho_n^*\} = \arg \min_{\mathcal{L}, \{\rho_n'\} \in \text{SAMPLE}(\{\rho_n\})} \ell(\mathcal{L}, \{\rho_n'\}), \quad (2)$$

173 for a chosen metric M . The choice of M is analysed in Section 6.

174 5.1 Sample with clustering

175 Meaningful abstractions arise when programs share underlying functionality or structure. To surface
 176 these, we cluster input programs into small groups that are likely to share reusable structure. Most
 177 modern language models cannot be prompted with the entire collection of input programs— even long
 178 context models cannot process the entirety of e.g the Linux kernel, and even if they could, it is not
 179 clear that such a strategy is the most efficient way of focusing the language model’s attention.

180 We consider clustering algorithms for discovering small groups of related code; we call these *tuples*.
 181 This extends REGAL [3], which clusters programs solving similar problems by assuming each
 182 program is paired with a natural language description of the problem it solves, and clustering
 183 embeddings of those descriptions. Since similar problems need not imply similar code structure, we
 184 instead prompt a model to summarize each code source and cluster by these summaries.

185 Each tuple is refactored in two stages: (1) if a library has already been built from previous tuples,
 186 we prompt the model with the current tuple and the accumulated library, asking it to identify which
 187 existing functions can be reused. This lets abstractions discovered earlier carry forward across the
 188 collection; (2) the retrieved functions and the tuple of programs are then provided as context to the
 189 model, which proposes a sample budget of K candidate refactorings.

190 5.2 Rank refactorings

191 Once sampled, all K candidate refactorings are passed through a *sort-and-filter* evaluation harness
 192 to select the refactoring that (1) scores the highest on refactor quality (under metric M) and (2)
 193 maintains (or improves) test accuracy compared to the original. If no such candidate exists, the
 194 original code is preserved, maintaining existing functionality.

195 New library functions in the selected refactor are saved into the LIBRARIAN library for potential use
 196 in downstream refactoring of other programs. We provide the full algorithm in Appendix A.

197 6 What Makes a Good Refactoring?

198 We evaluate candidate metrics along two complementary axes: (i) how they behave as optimization
 199 objectives when varying the sample budget K (Fig.3); and (ii) how well they align with human
 200 judgment of refactoring quality (Fig.4).

201 6.1 Objective function comparison

202 We run LIBRARIAN on 6 collections of CodeCon-
 203 tests using MDL, number of tokens, and cyclomatic
 204 complexity [20]² as objective functions M (Figure 3).
 205 While minimizing MDL also minimizes the other
 206 two objectives, the converse is not true. This suggests
 207 that MDL is a pareto-optimal loss among the three
 208 objectives in this experiment.

209 To confirm that the library does indeed expose shared
 210 abstractions, we calculate the average number of
 211 times that each library routine is used. Scaling the
 212 inference budget to $K = 8$ discovers better libraries,
 213 reusing each library function on average about 5
 214 times.

215 This scaling benefit holds for complex, real-world
 216 code as well, as demonstrated by monotonic com-
 217 pression improvements on the Transformers codebase
 218 (Figure 2). This reinforces our choice of MDL as a
 219 robust search objective. We prove that our best@k
 220 metric provides an unbiased estimate of this

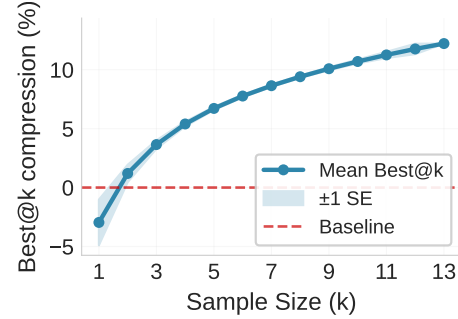


Figure 2: Best@k compression for LIBRARIAN on the Transformers domain. Increasing the sample budget (k) improves compression.

221 6.2 Human evaluation best aligns with the MDL objective function

222 We complement the quantitative analysis with a small-
 223 scale human study: for 19 CodeContests tuples of 3
 224 problems and a pair of refactorings for each, 4 hu-
 225 man judges pick which refactoring they prefer. Each
 226 refactoring comprises three re-written programs, plus
 227 the library generated by LIBRARIAN. Each pair of
 228 refactorings *passed all testcases*, and were chosen to
 229 include both the highest and lowest MDL refactor-
 230 ings. Human judges are authors on this paper, but
 231 were ‘blinded’ to not know relative objective function
 232 scores for the candidate refactorings.

233 An MDL objective aligns best with human prefer-
 234 ences (Figure 4), although number of tokens is also
 235 a reasonable proxy for human judgment. That token
 236 count and MDL yield different statistics furthermore
 237 implies that the lowest MDL program is not always the shortest program. LLM-as-judge (LLM-C)
 238 shows promising results, but currently underperforms MDL. Maintainability-Index³ and Cyclomatic-
 239 Complexity perform the worst on our test.

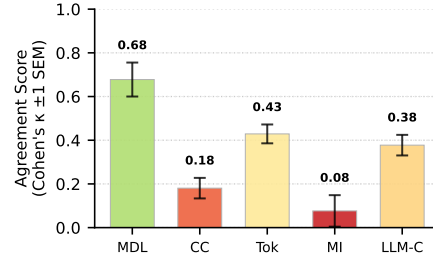


Figure 4: Human evaluation of different refactoring objectives. Judges compare pairs of refactorings that both pass all test cases. MDL aligns best with human preferences.

²Cyclomatic complexity (CC) is a classic metric from the software engineering community that analyses the number of paths through a program’s control flow, and unlike MDL is only loosely coupled with program size.

³Maintainability Index (MI) is a composite software engineering metric that combines lines of code, cyclomatic complexity, and Halstead volume into a single score. Higher MI values are intended to indicate easier-to-maintain code.

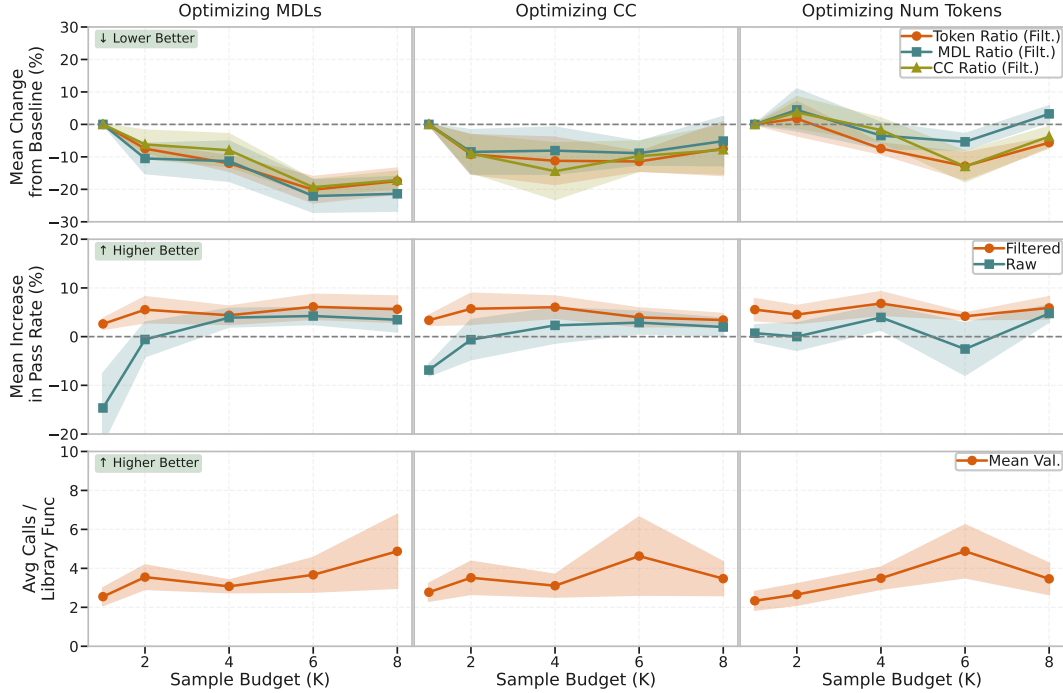


Figure 3: Comparing 3 different objective functions for refactoring (different columns) according to different downstream success metrics (different rows), as a function of refactoring budget (horizontal axes). The values are averaged over 6 collections of CodeContests problems. Row 1: Optimizing perplexity also incidentally optimizes cyclomatic complexity and token count, but that the converse is not true. Row 2: refactored programs pass more test cases, even more than the original code itself. Row 3: increasing the refactoring budget results in more reusable library subroutines (such subroutines are called more times on average). Filtered/Raw: Using/Not Using tests to filter samples.

Together, these experiments suggest that MDL is the most effective objective for guiding refactoring: it is Pareto-optimal among common complexity metrics, scales with inference budget to discover more reusable libraries, and best aligns with human judgments of quality. **We adopt M_{MDL} as the primary objective in the remainder of this paper.**

7 Experimental Setup

Grouping Programs into Collections To facilitate parallel application of LIBRARIAN and manage the dataset scale, we assume that semantically distant code sources will have minimal overlap in their optimal library functions. Therefore, our overall approach partitions the dataset into disjoint collections through clustering.

For **CodeContests**, these collections are constructed from an initial corpus of $\sim 9k$ problems with Python solutions: We first filter these code sources, removing those whose selected canonical solution is under 10 lines (minimal refactoring potential). For the remaining 4596 solutions we use a language model to generate textual descriptions of canonical solutions—emphasizing reusable components—which are embedded using OpenAI’s `text-embedding-ada-002`.

Agglomerative Clustering [22] is subsequently applied to these embeddings to partition the code sources into a predefined number of initial clusters, in our case 120. To create uniformly sized experimental units, we subsample each such cluster to form collections of 30 code sources. This collection size was empirically chosen because it balanced between the runtime of LIBRARIAN without limiting compression. We select 10 collections that we then use to evaluate our methods.

Code repositories are generated as disjoint collections through the generative process, which first samples project ideas then variations. We sample two code sources for each of the 10 collections. Repositories are setup as multi-package monorepositories. These monorepos are constructed by merging the dependencies and tests of the initial packages, and placing the source directories into the root directory of the monorepo.

For Transformers, since the number of models is on the lower end, we manually chose a set of popular LLM / VLM models and passed them to the agent in collections of 5 code sources.

REGAL Baselines. To evaluate the ability of our libraries to support reuse on new problems, we turn to the program synthesis tasks used in REGAL, where learned libraries are added to help the program synthesizer. We evaluate on the two domains published by the authors, Logo and Date. Because our clustering is inspired by REGAL but adds additional complexity, for fair comparison, we keep their setup the same and only augment the training using sample + MDL rerank procedure described in Section 5.1.

Code Contests. To evaluate LIBRARIAN on refactoring Code Contests we select 6 collections of 30 code sources (problems). In each collection we group the problems into tuples of size 3. We set the sample budget to be $K = 8$, since our ablations show that with larger K we discover better libraries. We use the MDL objective for rankings. The model used for sampling is OpenAI’s o4-mini [23]. To obtain MDL scores we use Qwen 2.5 7B Instruct [24] as a balance between quality, speed, and cost.

Code Agents on Transformers and Synthetic Repositories. To fairly evaluate performance on the task by state-of-the-art systems, we use coding agents that advertise long-context ability to reason about, write, and refactor code repositories. Specifically, we use Claude Code (Cl) [25] which uses the Claude 3.7 Sonnet model.

We test whether code agents can refactor collections of code sources autonomously, without human intervention. Refactoring repositories with code agents involves planning and iterative (re-)implementation and testing. Code agents are prompted to perform each of these steps, with feedback from the unit tests. Agents must run and repair unit tests autonomously. We run coding agents multiple times per collection, logging their progress in checklists stored in text files.

Due to context-length limitations in Qwen 2.5 7B, we instead use DeepSeek-V3 to compute MDL scores.

8 Results

In this section, we present the compression and correctness results with LIBRARIAN and agent baselines on MINICODE.

MINICODE-CodeContests. On CodeContests, LIBRARIAN achieves a high final pass rate of 90.67% and significantly improves correctness, with pass rates increasing by 6.33% compared to the original code sources (Table 5). The method yields substantial compression: the refactored code, including the new library, shows an MDL ratio of 0.53 (a 47% reduction in MDL relative to the original). On average, LIBRARIAN generates libraries containing approximately 11 functions. These functions demonstrate good reuse, being called by around 5.2 programs on average, although 38.03% of them are used only once within their specific collection context.

Code agents fail to achieve both high correctness and compression on MINICODE-CodeContests. Across collections, the Codex agent achieves an average MDL ratio of 0.83, but a pass rate of 74.16%, much lower than LIBRARIAN’s rate of 90.67%. Similarly, the Claude agent reaches a higher pass rate of 82.50%, which is still lower than LIBRARIAN, but an MDL ratio of 1.07 which is more complex than the original collection. We present the full results of the agents in Appendix E, along with results based on newer models such as Claude Sonnet 4 and codex-mini. With Claude Code and Sonnet 4,

Metric	Value
Pass Rate	90.67% $\pm$ 1.88
Pass Rate Improvement	6.33% $\pm$ 1.41
MDL Ratio	0.53 $\pm$ 0.03
Token Ratio	0.66 $\pm$ 0.04
Library Functions	10.30 $\pm$ 1.41
Avg Calls per Function	5.17 $\pm$ 1.08
% Single Use Functions	38.03% $\pm$ 4.88

Figure 5: Refactoring results for LIBRARIAN (w/ $K = 8$) averaged over 10 Code Contests collections.

agents achieve an MDL ratio of 0.77 and pass rate of 84.4%, outperforming codex-mini at an MDL ratio of 86.8 and pass rate of 82.0%.

MINICODE-Transformers. On Transformers refactoring task Claude Code achieved reasonable MDL ratio of 0.91 when using sampling budget of $K = 8$, while still passing the integration tests. The agent was able to extract repeated patterns such as MLP, Attention, Decoder classes, RoPE helper functions, etc. The main limitation of this was the computational cost of producing refactorings, since each refactoring took approximately 30 minutes to execute, limiting the number of samples we can reasonably do. This domain shows that sample-and-rerank method improves refactoring performance even on real world, larger scale repositories like Transformers.

MINICODE-Repositories. Coding agent results on synthetic repositories, generated by Sonnet 3.7, are given in Table 6. A state-of-the-art coding agent, Claude Code with Sonnet 3.7, was unable to achieve compression. The generated refactorings are larger than the original repositories, and often contain duplicate implementations due to incomplete refactorings. This indicates weaknesses in long-horizon agency for Sonnet 3.7. We provide examples in Appendix H.

REGAL Baseline Results. Beyond its success on our large-scale MINICODE benchmark, we sought to verify that our core sample-and-rerank method is general enough to benefit other state-of-the-art systems. To this end, we integrated our MDL-based reranking (using $K = 5$) into the REGAL framework for program synthesis. Despite the relative simplicity of the Logo and Date domains compared to our own, this augmentation yielded significant performance gains over the original REGAL baseline (Figure 7). This result demonstrates our method can be used to improve established methods in the classic library learning paradigm.

Collection		MDL %	Pass %
filesys_analyzer	orig	100	100
	claude	160	100
query_language	orig	100	100
	claude	140	100
knowledge_store	orig	100	100
	claude	140	100
vm_emulator	orig	100	100
	claude	160	100
finance_tracker	orig	100	100
	claude	180	60
in_memory_db	orig	100	100
	claude	150	100
text_editor	orig	100	100
	claude	120	100

Figure 6: Correctness (Pass %) and compression (MDL ratio as %) of original and refactored code sources for agent baselines on repository collections in MINICODE.

9 Conclusion

We introduce a new benchmark MINICODE and method LIBRARIAN for compressing code sources through reusable abstractions. We highlight the challenges that modern models face when producing modular and maintainable code, then present an effective method for using LLMs to do this task in small-scale programs. By framing refactoring as both a design and compression task, our work opens new directions for building more general and scalable code understanding and generation systems. In particular, the structure of MINICODE lends itself well to reinforcement learning, where training would entail synthesizing collections of repositories to refactor then computing rewards based on compression.

Limitations Limitations of this work include the evaluation of synthetic repositories, the poor performance of code agents on refactoring, and the fact that compression may not be correlated with reuse. While our experiments on downstream programming problems partly address the question of reuse, investigating how to measure and encourage reuse for large-scale multi-repo library creation remains an open problem.

Figure 7: Solving downstream program synthesis tasks using learned libraries

Dataset	Model	Pass Rate
Logo	REGAL (gpt-3.5-turbo)	49.3% ± 1.1
	LIBRARIAN (3.5-turbo)	69.9% ± 0.9
Date	REGAL (gpt-3.5-turbo)	90.2% ± 0.5
	LIBRARIAN (3.5-turbo)	94.7% ± 0.7

References

- [1] Shane Tews. Inside tech’s \$2 trillion technical debt | american enterprise institute - aei.
- [2] Catherine Wong, Kevin Ellis, Joshua B. Tenenbaum, and Jacob Andreas. Leveraging language to learn program abstractions and search heuristics. In *International Conference on Machine Learning*, 2021.
- [3] Elias Stengel-Eskin, Archiki Prasad, and Mohit Bansal. Regal: refactoring programs to discover generalizable abstractions. In *Proceedings of the 41st International Conference on Machine Learning, ICML’24*. JMLR.org, 2024.
- [4] Kevin Ellis, Catherine Wong, Maxwell Nye, Mathias Sablé-Meyer, Lucas Morales, Luke Hewitt, Luc Cary, Armando Solar-Lezama, and Joshua B. Tenenbaum. Dreamcoder: Bootstrapping inductive program synthesis with wake-sleep library learning. In *PLDI*, 2021.
- [5] Matthew Bowers, Theo X. Olausson, Lionel Wong, Gabriel Grand, Joshua B. Tenenbaum, Kevin Ellis, and Armando Solar-Lezama. Top-down synthesis for library learning. *Proc. ACM Program. Lang.*, 7(POPL), January 2023.
- [6] Eyal Dechter, Jon Malmaud, Ryan P. Adams, and Joshua B. Tenenbaum. Bootstrap learning via modular concept discovery. In *IJCAI*, 2013.
- [7] Gabriel Grand, Li Siang Wong, Matthew Bowers, Theo X. Olausson, Muxin Liu, Joshua B. Tenenbaum, and Jacob Andreas. Lilo: Learning interpretable libraries by compressing and documenting code. *ArXiv*, abs/2310.19791, 2023.
- [8] Percy Liang, Michael I Jordan, and Dan Klein. Learning programs: A hierarchical bayesian approach. In *ICML*, volume 10, pages 639–646, 2010.
- [9] Carlos E Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Pei, Ofir Press, and Karthik R Narasimhan. SWE-bench: Can language models resolve real-world github issues? In *The Twelfth International Conference on Learning Representations*, 2024.
- [10] Wenting Zhao, Nan Jiang, Celine Lee, Justin T Chiu, Claire Cardie, Matthias Gallé, and Alexander M Rush. Commit0: Library generation from scratch. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [11] Dhruv Gautam, Spandan Garg, Jinu Jang, Neel Sundaresan, and Roshanak Zilouchian Moghaddam. Refactorbench: Evaluating stateful reasoning in language agents through code. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [12] Zhiruo Wang, Daniel Fried, and Graham Neubig. Trove: Inducing verifiable and efficient toolboxes for solving programmatic tasks, 2024.
- [13] Siddhant Waghjale, Vishruth Veerendranath, Zora Zhiruo Wang, and Daniel Fried. Ecco: Can we improve model-generated code efficiency without sacrificing functional correctness? *arXiv preprint arXiv:2407.14044*, 2024.
- [14] Anne Ouyang, Simon Guo, Simran Arora, Alex L. Zhang, William Hu, Christopher Ré, and Azalia Mirhoseini. Kernelbench: Can llms write efficient gpu kernels?, 2025.
- [15] Eric Schkufza, Rahul Sharma, and Alex Aiken. Stochastic superoptimization. In *Proceedings of the Eighteenth International Conference on Architectural Support for Programming Languages and Operating Systems*, ASPLOS ’13, page 305–316, New York, NY, USA, 2013. Association for Computing Machinery.
- [16] Oleksandr Polozov and Sumit Gulwani. Flashmeta: A framework for inductive program synthesis. *ACM SIGPLAN Notices*, 50(10):107–126, 2015.
- [17] David Cao, Rose Kunkel, Chandrakana Nandi, Max Willsey, Zachary Tatlock, and Nadia Polikarpova. Babble: Learning better abstractions with e-graphs and anti-unification. *POPL*, 2023.
- [18] Percy Liang, Michael I. Jordan, and Dan Klein. Learning programs: A hierarchical bayesian approach. In *ICML*, 2010.
- [19] Ray J Solomonoff. A formal theory of inductive inference. *Information and control*, 7(1):1–22, 1964.
- [20] T.J. McCabe. A complexity measure. *IEEE Transactions on Software Engineering*, SE-2(4):308–320, 1976.
- [21] Yujia Li, David Choi, Junyoung Chung, Nate Kushman, Julian Schrittwieser, Rémi Leblond, Tom Eccles, James Keeling, Felix Gimeno, Agustin Dal Lago, Thomas Hubert, Peter Choy, Cyprien de Masson d’Autume, Igor Babuschkin, Xinyun Chen, Po-Sen Huang, Johannes Welbl, Sven Gowal, Alexey Cherepanov, James Molloy, Daniel Mankowitz, Esme Sutherland Robson, Pushmeet Kohli, Nando de Freitas, Koray Kavukcuoglu, and Oriol Vinyals. Competition-level code generation with alphacode. *arXiv preprint arXiv:2203.07814*, 2022.
- [22] Joe H Ward Jr. Hierarchical grouping to optimize an objective function. *Journal of the American statistical association*, 58(301):236–244, 1963.

- 419 [23] OpenAI. Introducing o3 and o4 mini. [https://openai.com/index/](https://openai.com/index/introducing-o3-and-o4-mini/)
420 [introducing-o3-and-o4-mini/](https://openai.com/index/introducing-o3-and-o4-mini/), 2024. Accessed: 2025-05-13.
- 421 [24] Qwen, :, An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan Li,
422 Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang,
423 Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin Yang, Le Yu, Mei Li,
424 Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li, Tianyi Tang, Tingyu Xia, Xingzhang
425 Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan, Yuqiong Liu, Zeyu Cui, Zhenru Zhang,
426 and Zihan Qiu. Qwen2.5 technical report, 2025.
- 427 [25] Anthropic. Claude code: An agentic coding tool that lives in your terminal. [https://github.com/](https://github.com/anthropics/claude-code)
428 [anthropics/claude-code](https://github.com/anthropics/claude-code), 2025. An agentic coding tool that lives in your terminal, understands your
429 codebase, and helps you code faster by executing routine tasks, explaining complex code, and handling git
430 workflows - all through natural language commands.

Algorithm 1 Refactoring Specialized Programs into a Joint Library

Require: Set of independent, specialized programs $P_{initial} = \{\rho_1, \rho_2, \dots, \rho_n\}$ **Require:** Sample Budget K **Ensure:** Joint library $\mathcal{L}_{final}$ and set of refactored programs P_{final}

```

1:  $\mathcal{C} \leftarrow \text{Cluster}(P_{initial})$ 
2:  $\mathcal{L}_{final} \leftarrow \emptyset, P_{final} \leftarrow \emptyset$ 
3: for all cluster  $c \in \mathcal{C}$  do                                 $\triangleright$  Each cluster independently
4:    $T_C \leftarrow \text{GroupIntoTuples}(c)$                          $\triangleright$  Get tuples for each cluster
5:   for all tuples  $\tau \in T_C$  do
6:      $\{f_{retrieved}\} \leftarrow \text{RetrieveRelevantFromLibrary}(\mathcal{L}, \tau)$ 
7:      $S \leftarrow \emptyset$ 
8:     for  $i = 1$  to  $K$  do                                        $\triangleright$  Sample  $k$  times
9:        $(\{f_{new,i}\}, \{\rho'_i\}) \leftarrow \text{SAMPLE}(f_{retrieved}, \tau)$ 
10:       $S \leftarrow S \cup \{(\{f_{new,i}\}, \{\rho'_i\})\}$ 
11:    end for
12:     $(f_{best}, \{\rho'_{best}\}) \leftarrow \text{RerankAndSelectBest}(S, \ell(\cdot))$   $\triangleright$  Rerank using objective
13:     $\mathcal{L}_{final} \cup \{f_{best}\}$ 
14:     $P_{final} \cup \{\rho'_{best}\}$ 
15:  end for
16: end for
17: return  $\mathcal{L}_{final}, P_{final}$ 

```

432 The instruction provided to human evaluators is as follows:

```
1
2
3 ## 1. Materials Provided
4
5 You will be given a set of files for each example case:
6
7 * **'original_programs.py':** This file contains a set of 3 distinct Python programs, each presented
8   with its corresponding problem description/query. This represents the "before" state.
9 * **'v1.py':** This file presents the first refactoring approach. It includes:
10   * The 3 refactored versions of the original programs.
11   * A "library" section (e.g., 'codebank.py' or inline) containing helper functions. These helper
12     functions might be retrieved from an existing common library or newly created during this
13     refactoring.
14   * Either the retrieved or the new helper function sections may be _empty_, in case no programs
15     existed in the codebank at the time or if no need helper functions were created by the LLM.
16 * **'v2.py':** This file presents the second, alternative refactoring approach. Similar to '
17   refactoring_v1.py', it includes:
18   * The 3 refactored versions of the original programs (using a different strategy than v1).
19   * A "library" section with its own set of helper functions.
20
21 **NOTE**: both refactorings had accuracy at least as good as the original programs.
22
23 ## 2. Your Task
24
25 Your primary task is to:
26
27 1. **Review** the 'original_programs.py' to understand the initial code and the problems being solved.
28 2. **Analyze** both 'refactoring_v1.py' and 'refactoring_v2.py'. Pay close attention to how the
29   original programs have been restructured and what functionalities have been extracted into their
30   respective libraries.
31 3. **Decide** which refactoring (Version 1 or Version 2) you believe is "better,"** based on the
32   evaluation criteria provided below (or your own criteria!).
33
34 ## 3. Evaluation Criteria: What to Consider for Your Choice
35
36 When comparing 'refactoring_v1.py' and 'refactoring_v2.py', please *consider* the following aspects to
37 inform your choice. The "better" refactoring should ideally excel in these areas:
38
39 > Most importantly, make sure that the extracted functions are **actually reusable and not too specific
40   ** If the main programs are short, the refactoring is not immediately "better"! Try to think
41   whether the extracted functions could actually be used in a different program down the line.
42
43 * **Reusability of Helper Functions:**
44   * **Generality:** Are the new helper functions general-purpose and potentially useful for *other,
45     different* programs and problems beyond the three presented?
46   * **Reuse:** How much were existing helper functions reused?
47   * **Specificity:** Are the functions too specialized to the current set of problems, limiting their
48     broader applicability? _Avoid functions that are essentially just the original program broken out
49     into a "helper."_
50   * Composability
51 * **Maintainability:**
52   * Readability & Understandability
53   * Ease of Modification
54   * Separation of Concerns
55
56 ## 4. What NOT to Focus On:
57
58 * **Comments:** Please disregard the presence or absence of comments in the code for this evaluation.
59   These are superficially generated by LLMs in some occasions and could be added manually after with
60   a single pass.
61 * **Minor Stylistic Differences:** Do not focus on trivial differences in variable naming or formatting
62   , unless they significantly impact readability or understanding.
63
64 ## 5. How to Provide Your Feedback
65
66 For each example case, please provide:
67
68 1. **Your Preferred Version:** (e.g., "Version 1" or "Version 2")
```

Listing 1: Human Evaluation Instruction

433 B Best@k Compression is a U-Statistic

434 We wish to estimate the expected compression ratio achieved by our *sample + rerank* method, which
435 samples k candidate refactorings, discards any that do not pass the tests, and selects the one with the
436 lowest score (total log-prob).

437 **Background on U-Statistics.** Let $Z_1, \dots, Z_n \stackrel{\text{i.i.d.}}{\sim} F$. For a symmetric function $h : \mathcal{Z}^k \rightarrow \mathbb{R}$, the
 438 *U-statistic of order k* is defined as

$$U_n = \binom{n}{k}^{-1} \sum_{1 \leq i_1 < \dots < i_k \leq n} h(Z_{i_1}, \dots, Z_{i_k}). \quad (4)$$

439 By construction,

$$\mathbb{E}[U_n] = \mathbb{E}[h(Z_1, \dots, Z_k)], \quad (5)$$

440 so U_n is an unbiased estimator of the population quantity $\theta = \mathbb{E}[h(Z_1, \dots, Z_k)]$.

441 **Application to Best@ k Compression.** Let each valid refactoring be a pair $Z = (S, C)$, where S is
 442 the score and C is the compression ratio. Define the symmetric function

$$h_k(z_1, \dots, z_k) = C_{j^*}, \quad j^* = \arg \min_{1 \leq j \leq k} S_j, \quad (6)$$

443 the compression ratio of the lowest-score refactoring among k draws. The population target is then

$$\theta_k = \mathbb{E}[h_k(Z_1, \dots, Z_k)]. \quad (7)$$

444 Given n valid samples, our estimator is

$$\hat{\theta}_k = \binom{n}{k}^{-1} \sum_{1 \leq i_1 < \dots < i_k \leq n} h_k(Z_{i_1}, \dots, Z_{i_k}). \quad (8)$$

445 **Proposition.** $\hat{\theta}_k$ is a U-statistic of order k with function h_k , and hence an unbiased estimator of θ_k .

446 **Proof.** (1) *Symmetry of the function h_k .* h_k selects the compression associated with the lowest score
 447 among its k arguments. Permuting the inputs does not affect this outcome (ties can be resolved with
 448 a fixed, permutation-invariant rule). Thus h_k is symmetric.

449 (2) *U-statistic form.* By definition, a U-statistic of order k with kernel h_k is

$$U_n = \binom{n}{k}^{-1} \sum_{1 \leq i_1 < \dots < i_k \leq n} h_k(Z_{i_1}, \dots, Z_{i_k}),$$

450 which matches $\hat{\theta}_k$ exactly.

451 Therefore, $\hat{\theta}_k$ is a U-statistic of order k . By the unbiasedness property of U-statistics,

$$\mathbb{E}[\hat{\theta}_k] = \theta_k.$$

452 □

453 Thus, our reported *best@ k compression curves* provide unbiased estimates of the expected perfor-
 454 mance of the *sample + rerank* method.

455 C Clustering Analysis: CodeContests

456 We analyze the coherence of the clusters underlying collections in MINICODE-**CodeContests**. In
 457 particular, we compare clustering based on o4-mini generated code source descriptions against
 458 task descriptions. Since task descriptions in competition coding problems are designed to hide
 459 the algorithmic approach needed to solve problem, we expect that clusters based on code source
 460 descriptions are more coherent. We use Normalized Tag Instance Entropy and Herfindahl-Hirschman
 461 Index to evaluate clusterings. Figure 8 shows our clustering approach yields more thematically
 462 coherent clusters, evidenced by achieving lower entropy and higher HHI values across the entire
 463 tested range of N . We provide definitions of our measures below.

464 C.1 Collection Coherence Measures

465 We use two measures to evaluate the thematic coherence of collections: Good collections should
 466 group code sources with a (1) concentrated and (2) identifiable set of shared *conceptual tags*, which
 467 for CodeContests are provided as ground truth (`trees`, `graphs`, etc.).

468 We provide the full definitions of the collection coherence measures below.

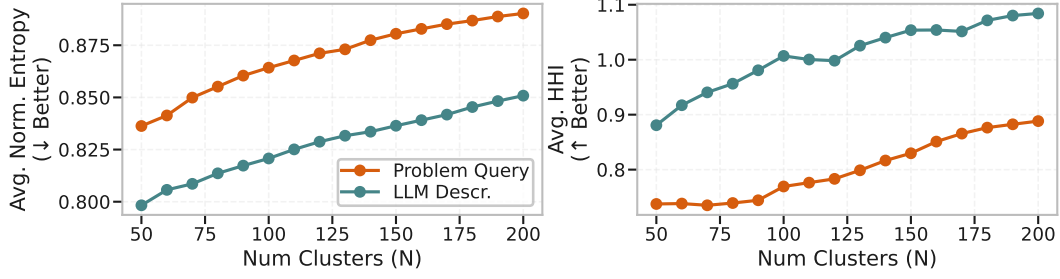


Figure 8: Clustering analysis of 4,596 Code Contest problems, comparing the thematic coherence of clusters formed using our proposed method versus REGAL-style clustering.

469 **Normalized Tag Instance Entropy:** This measures the concentration of tag *instances* within a
 470 collection C . Let p_i be the proportion of the i -th unique tag type among all tag instances in C , and
 471 D_C be the number of distinct tag types in C . If $D_C > 1$, the normalized entropy H_N is defined as:

$$H_N = -\frac{\sum_{i=1}^{D_C} p_i \log_2 p_i}{\log_2 D_C} \quad (9)$$

472 If $D_C \leq 1$, then $H_N = 0$. Lower H_N (closer to 0) indicates higher thematic purity, meaning fewer
 473 tag types dominate the bulk of tag mentions.

474 **Herfindahl-Hirschman Index (HHI) for Problem Presence:** This measures tag concentration
 475 across distinct *problems* in a cluster C . Let s_t be the proportion of problems in C that include tag t (a
 476 problem contributes to s_t if t is one of its unique tags). A higher HHI signifies that the problems are
 477 collectively characterized by a smaller, more focused set of tags.

$$\text{HHI} = \sum_{t \in \text{Tags}(C)} s_t^2 \quad (10)$$

478 where $\text{Tags}(C)$ represents the set of unique tags present in cluster C .

479 D Benchmark Comparison

480 We compare our benchmark, MINICODE, to similar benchmarks in Table 2. We define creativity and
 481 design as the need to explore diverse solutions in order to find the best solution possible. For example,
 482 optimizing for program correctness alone does not require exploring a large solutions space, whereas
 483 optimizing a program for speed would. In the case of compressing large code sources, we must
 484 explore the large space of shared abstractions afforded by libraries in order to maximize compression.

Table 2: Comparison of Code Benchmarks

Benchmark	Creativity/Design	Scale
SWE-bench [9]	Low	Repository
Commit-0 [10]	Medium	Repository
RefactorBench [11]	Low	File
ECCO [13]	High	Function
KernelBench [14]	High	Function
MINICODE(Ours)	High	Multi-repository

485 E Full MINICODE Results

486 We present the full agent scores for the CodeContests split in Table 3. The results are given both for
 487 each cluster of code sources, as well as averaged across clusters.

Cluster	Agent	Tokens	CC	Pass %	MDL	MDL %
0	original	9088	95	80.3	11745.85	100.0
	sonnet 3.7	18114	176	87.0	15005.18	127.7
	sonnet 4	11121	138	80.3	9901.53	84.3
	codex-mini	9321	95	80.3	9990.74	85.1
1	original	12531	255	89.7	13431.86	100.0
	sonnet 3.7	10470	239	96.7	8933.65	66.5
	sonnet 4	11325	298	96.7	8214.42	61.2
	codex-mini	12762	255	89.7	11798.73	87.8
2	original	14087	376	89.0	15012.77	100.0
	sonnet 3.7	17345	429	91.3	13145.02	87.6
	sonnet 4	14270	356	93.0	10522.66	70.1
	codex-mini	14318	376	89.0	13273.81	88.4
3	original	14261	246	90.3	13348.82	100.0
	sonnet 3.7	20749	241	97.7	15859.02	118.8
	sonnet 4	13433	197	80.7	11937.04	89.4
	codex-mini	14495	246	90.3	11616.41	87.0
4	original	17693	336	80.7	14665.16	100.0
	sonnet 3.7	29860	358	100.0	20666.52	141.0
	sonnet 4	18684	352	82.0	12801.21	87.3
	codex-mini	17923	336	80.7	12902.09	88.0
5	original	12588	286	92.0	12790.11	100.0
	sonnet 3.7	10580	128	99.3	8435.12	65.9
	sonnet 4	10416	155	99.3	9167.85	71.7
	codex-mini	12819	286	92.0	11086.19	86.7
6	original	11020	131	54.3	13540.41	100.0
	sonnet 3.7	21747	502	88.0	19446.07	143.6
	sonnet 4	11177	143	57.3	10492.00	77.5
	codex-mini	11251	131	54.3	11651.65	86.1
7	original	12301	180	80.0	12393.73	100.0
	sonnet 3.7	16390	166	91.0	13371.59	107.9
	sonnet 4	11625	150	85.7	9304.25	75.1
	codex-mini	12534	180	80.0	10549.04	85.1
Avg	original	12946	238	82.0	13366.09	100.0
	sonnet 3.7	18157	280	93.9	14357.77	107.4
	sonnet 4	12756	224	84.4	10292.62	77.1
	codex-mini	13178	238	82.0	11608.58	86.8

Table 3: Comparison of the pass rate and compression metrics of the original code sources, Claude Sonnet 4 and codex-mini refactorings across CodeContests clusters.

488 We present the full repository-level results of MINICODE-repository in Tables 4 and 5, with o4-mini
489 and Claude Sonnet 3.7 and 4.

Collection	Agent	LLoC	CC	MDL ratio	Pass rate
datapipe	original	474	116	1.0	1.0
	Cl-Cl	1084	341	9.1	1.0
	Cl-Cx	645	174	2.2	1.0
state_machine	original	619	222	1.0	1.0
	Cl-Cl	2271	735	6.3	0.9
	Cl-Cx	686	227	1.9	1.0
config_schema	original	949	284	1.0	1.0
	Cl-Cl	922	339	2.5	failed
	Cl-Cx	626	207	2.1	1.0
cli_tools	original	875	300	1.0	1.0
	Cl-Cl	2925	909	4.0	1.0
	Cl-Cx	5848	2378	3.1	failed
cli_form	original	352	174	1.0	1.0
	Cl-Cl	855	342	3.8	1.0
	Cl-Cx	1366	548	2.9	1.0

Table 4: Full results on MINICODE-repositories small, using Codex with o4-mini and Claude Code with Claude Sonnet 3.7.

Collection	Agent	LLoC	CC	MDL %	Pass %
command_line_task_manager	original	7964	2921	100	100
	sonnet 3.7	10387	3899	130	100
	sonnet 4	9515	3946	125	100
concurrent_task_scheduler	original	10350	3935	100	100
	sonnet 3.7	18653	7165	190	80
	sonnet 4	11491	4762	122	80
file_system_analyzer	original	3911	1338	100	100
	sonnet 3.7	6495	2218	160	100
	sonnet 4	5221	1793	130	100
in_memory_database	original	4565	1671	100	100
	sonnet 3.7	6667	2556	150	100
	sonnet 4	5617	2356	131	100
incremental_backup_system	original	4587	1482	100	100
	sonnet 3.7	5365	1788	130	failed
	sonnet 4	6620	2168	152	60
personal_finance_tracker	original	4306	1482	100	100
	sonnet 3.7	7443	2512	180	63
	sonnet 4	9212	3335	215	63
personal_knowledge_management	original	5341	1832	100	100
	sonnet 3.7	6357	2364	140	100
	sonnet 4	5575	2597	131	80
query_language_interpreter	original	5181	2105	100	100
	sonnet 3.7	7193	2966	140	100
	sonnet 4	7490	3076	142	100
text_editor	original	4324	1323	100	100
	sonnet 3.7	5792	1822	140	100
	sonnet 4	6017	2308	148	100
virtual_machine_emulator	original	6841	2283	100	100
	sonnet 3.7	10108	3580	160	100
	sonnet 4	9220	3303	137	100

Table 5: Full results on MINICODE-repositories large, comparing the original code sources with Claude Sonnet 3.7 and Sonnet 4.

490 F Data Generating Prompts

491 We provide the prompts for Librarian here.

```

1 I need ideas for Python libraries that can be implemented by language models. These libraries should:
2
3 1. Be implementable using only Python's standard library - no external dependencies
4 2. Have enough complexity to demonstrate sophisticated code design (10-20 functions/methods)
5 3. Include room for interpretation, so that different implementations can be unique while sharing core
   functionality
6 4. Have clear, practical utility that solves a real programming need
7 5. Be realistically implementable by an intelligent language model
8 6. Be testable with pytest
9 7. Include opportunities for different design approaches (functional vs OOP, etc.)
10
11 For each library, provide a description that outlines:
12 - The problem domain and core purpose
13 - Key required functionality (without being too prescriptive about implementation details)
14 - Potential use cases that demonstrate practical applications
15 - Suggested extension points where implementers could add their creative spin
16
17 Please generate several proposals in markdown, following this format:
18
19 ```file:<library_name>/DESCRIPTION.md
20 # <Library Name>
21
22 ## Purpose and Motivation
23 <3-5 sentences on what problem this library solves and why it's useful>
24
25 ## Core Functionality
26 <Description of 4-6 high-level key features/capabilities without specifying exact implementation>
27 ```
28
29 Be creative! Focus on domains where standard Python libraries provide enough building blocks but where
   a well-designed abstraction layer would add significant value.

```

Listing 2: Prompt to generate library descriptions

```

1
2 Consider a code repository designed to support the following task description:
3 ```
4 {task_content}
5 ```
6
7 Please list a couple dozen features that would be useful for the repository. Include the specified
   features as well as several others.
8
9 List the suggested feature names and descriptions in the following format:
10
11 1: <feature_1_name>: <feature_1_description>
12 2: <feature_2_name>: <feature_2_description>
13 3: <feature_3_name>: <feature_3_description>
14 ...
15 30: <feature_30_name>: <feature_30_description>
16 """
17
18 persona_prompt_template = """Consider the following features of a code repository:
19 {listed_features}
20
21 Think about several possibilities for what kind of person might use this code repository and what they
   might use it for. Please write several brief descriptions for the code repository in first person,
   formatted in markdown as follows:
22 ```file:{library_name}/<persona_name>/TASK.md
23 # The Task
24
25 I am a <...> I want to be able to <...> This code repository <...>
26
27 # The Requirements
28
29 * <function_name>: <feature description>
30 * ...
31 ```
32
33 Be creative! Write the task description in the style of the proposed persona. Be as exhaustive as
   possible in including the listed features in the task description's requirements.

```

Listing 3: Prompt to generate potential uses and features given a library description

```

1
2
3 I need you to implement a Python solution and COMPREHENSIVE suite of tests based on the following task
  files.
4 Your code must pass the tests provided.
5
6 {task_content}
7
8 CRITICAL FORMATTING INSTRUCTIONS:
9 1. You MUST format ALL code files exactly as shown below - no exceptions
10 2. Start each file with the markdown codeblock marker, followed by "file:" and the relative path
11 3. End each file with the closing markdown codeblock marker
12 4. Do not use any other format or markdown variations
13 5. For test files, do not put them in a subdirectory-- keep them in the outermost level.
14
15 For each source code file:
16
17 ``file:<relative_file_path>
18 <file_content>
19 ``
20
21 For each test file:
22
23 ``file:<relative_file_path starting with test_>
24 <test_file_content>
25 ``
26
27 IMPORTANT:
28 - The opening format must be exactly: ``file:path/to/file.py
29 - Do not add language indicators like ``python
30 - Do not add explanations between files
31 - Each file must be contained within its own codeblock with the precise format shown above
32 - The system parsing your response requires this exact format to function properly
33
34 EXAMPLE OUTPUT FORMAT:
35 ``file:mymodule/mymodule.py
36 def example_function():
37     return "This is a sample function"
38 ``
39
40 ``file:test_utils.py
41 import pytest
42 from mymodule.mymodule import example_function
43
44 def test_example_function():
45     assert example_function() == "This is a sample function"
46 ``
47
48 Begin your implementation now, following these formatting rules precisely.

```

Listing 4: Prompt to generate initial implementation

```

1
2 I need you to fix the implementation of the following code that is failing tests.
3
4 # Current Implementation:
5 {src_code_content}
6
7 # Test Files:
8 {test_content}
9
10 # Failed Tests:
11 {failed_test_details}
12
13 # Test Output (if available):
14 {test_output}
15
16 Please carefully analyze the errors and test failures. Pay special attention to:
17 1. The exact assertion failures or error messages
18 2. What the tests expect vs. what your current implementation provides
19 3. Any edge cases or special conditions you might have missed
20
21 Your task is to fix the implementation to make all tests pass. For each file that needs to be modified,
22 provide the content in the following format:
23
24 ““file:<relative_file_path>
25 <file_content>
26 ““
27
28 Where <relative_file_path> is the relative path to the file and <file_content> is the updated content
29 of the file.
30 Focus on fixing the specific issues identified in the errors and failed tests while maintaining the
31 overall structure of the code.
32
33 IMPORTANT: Make targeted changes to address the specific failing test cases. Make sure your
34 implementation passes all test cases,
35 including any edge cases or special conditions mentioned in the tests. Be sure to output code in the
36 specified format.

```

Listing 5: Prompt to fix code implementation, given pytest output.

492 G Refactoring examples of LIBRARIAN on Code Contests

493 G.1 Example 1

494 In code snippets 7, 6, 9, 8 one example of 2 refactoring versions. Specifically, the versions are
495 both passing at least as many test cases as the original and they have the biggest difference in
496 MDL among all the sample refactorings for that tuple. Sample + rerank filtering selected refac-
497 toring V2. You can observe that refactoring V1 introduces some problem specific functions like
498 `build_max_beauty_perm()`, while refactoring V2 sticks to more generally useful functions.

```
1 #==== NEW HELPER FUNCTIONS ====
2 def compute_full_mask(i):
3     """Return mask of all 1s of the bit-length of i."""
4     return (1 << i.bit_length()) - 1
5
6 def build_max_beauty_perm(n):
7     """Build permutation of 0..n maximizing sum of i^p[i]."""
8     ans = [0] * (n + 1)
9     used = set()
10    for i in range(n, -1, -1):
11        if i in used:
12            continue
13        mask = compute_full_mask(i)
14        j = i ^ mask
15        ans[i], ans[j] = j, i
16        used.add(i)
17        used.add(j)
18    beauty = sum(i ^ ans[i] for i in range(n + 1))
19    return ans, beauty
20
21 def solve_xor_sum(u, v):
22     """
23     Find shortest array whose xor is u and sum is v.
24     Return list or None if impossible.
25     """
26     if u > v or (v - u) % 2:
27         return None
28     if u == v:
29         return [] if u == 0 else [u]
30     x = (v - u) // 2
31     # try two elements
32     if (u + x) ^ x == u:
33         return [u + x, x]
34     # fallback to three elements
35     return [u, x, x]
36
37 def build_trie(keys):
38     """
39     Build a binary trie with counts for 30-bit numbers.
40     Each node: [left_index, right_index, count].
41     """
42     tree = [[0, 0, 0]]
43     for x in keys:
44         now = 0
45         tree[now][2] += 1
46         for i in range(29, -1, -1):
47             b = (x >> i) & 1
48             if tree[now][b] == 0:
49                 tree[now][b] = len(tree)
50                 tree.append([0, 0, 0])
51             now = tree[now][b]
52             tree[now][2] += 1
53     return tree
54
55 def trie_pop_min_xor(tree, x):
56     """
57     Pop one key from trie to minimize x^key and return that minimal xor.
58     Decrements counts along the path.
59     """
60     now = 0
61     res = 0
62     for i in range(29, -1, -1):
63         b = (x >> i) & 1
64         nxt = tree[now][b]
65         if nxt and tree[nxt][2] > 0:
66             now = nxt
67         else:
68             now = tree[now][b ^ 1]
69             res |= (1 << i)
70             tree[now][2] -= 1
71     return res
```

Listing 6: Version 1, New Helpers

```

1
2 # ##### PROGRAM: node_16:cc_python_16 #####
3
4 from codebank import *
5
6 def main():
7     import sys
8     data = sys.stdin.readline()
9     if not data:
10         return
11     n = int(data)
12     perm, beauty = build_max_beauty_perm(n)
13     print(beauty)
14     print(*perm)
15
16 if __name__ == "__main__":
17     main()
18
19 # ##### PROGRAM: node_19:cc_python_19 #####
20
21 from codebank import *
22
23 def main():
24     import sys
25     data = sys.stdin.readline()
26     n = int(data())
27     A = list(map(int, data().split()))
28     P = list(map(int, data().split()))
29     trie = build_trie(P)
30     O = [trie_pop_min_xor(trie, a) for a in A]
31     print(*O)
32
33 if __name__ == "__main__":
34     main()
35
36 # ##### PROGRAM: node_25:cc_python_25 #####
37
38 from codebank import *
39
40 def main():
41     import sys
42     u, v = map(int, sys.stdin.readline().split())
43     res = solve_xor_sum(u, v)
44     if res is None:
45         print(-1)
46     else:
47         print(len(res))
48         if res:
49             print(*res)
50
51 if __name__ == "__main__":
52     main()

```

Listing 7: Version 1, Refactored Programs

```

1 # ==== NEW HELPER FUNCTIONS ====
2 def compute_complement(i):
3     return i ^ ((1 << i.bit_length()) - 1)
4
5 def trie_add(trie, x, max_bit):
6     trie[0][2] += 1
7     now = 0
8     for i in range(max_bit, -1, -1):
9         bit = (x >> i) & 1
10         if trie[now][bit] == 0:
11             trie[now][bit] = len(trie)
12             trie.append([0, 0, 0])
13             now = trie[now][bit]
14             trie[now][2] += 1
15
16 def trie_find_min_xor(trie, x, max_bit):
17     now = 0
18     ans = 0
19     for i in range(max_bit, -1, -1):
20         bit = (x >> i) & 1
21         if trie[now][bit] and trie[trie[now][bit]][2] > 0:
22             now = trie[now][bit]
23         else:
24             now = trie[now][bit ^ 1]
25             ans |= (1 << i)
26         trie[now][2] -= 1
27     return ans

```

Listing 8: Version 2, New Helpers

```

1  ##### PROGRAM: node_16:cc_python_16 #####
2
3  from codebank import *
4
5  def main():
6      import sys
7      input = sys.stdin.readline
8      n = int(input())
9      ans = [-1] * (n + 1)
10     for i in range(n, -1, -1):
11         if ans[i] == -1:
12             z = compute_complement(i)
13             ans[i] = z
14             ans[z] = i
15     m = sum(i ^ ans[i] for i in range(n + 1))
16     print(m)
17     print(*ans)
18
19 if __name__ == "__main__":
20     main()
21
22 ##### PROGRAM: node_19:cc_python_19 #####
23
24 from codebank import *
25
26 def main():
27     import sys
28     input = sys.stdin.readline
29     n = int(input())
30     A = list(map(int, input().split()))
31     P = list(map(int, input().split()))
32     max_bit = max(max(A, default=0), max(P, default=0)).bit_length() - 1
33     trie = [[0, 0, 0]]
34     for x in P:
35         trie_add(trie, x, max_bit)
36     res = [trie_find_min_xor(trie, x, max_bit) for x in A]
37     print(*res)
38
39 if __name__ == "__main__":
40     main()
41
42 ##### PROGRAM: node_25:cc_python_25 #####
43
44 from codebank import *
45
46 def main():
47     u, v = map(int, input().split())
48     if u > v or ((v - u) & 1):
49         print(-1)
50     elif u == 0 and v == 0:
51         print(0)
52     elif u == v:
53         print(1)
54         print(u)
55     else:
56         w = (v - u) // 2
57         if (w & u) == 0:
58             d = u + w
59             print(2)
60             print(d, w)
61         else:
62             print(3)
63             print(u, w, w)
64
65 if __name__ == "__main__":
66     main()

```

Listing 9: Version 2, Refactored Programs

499 G.2 Example 2

500 In code snippets 11, 10, 13, 12 is another example of 2 refactorings where V1 was better according to
501 LIBRARIAN. We can observe that V2 creates helper functions that are overly specific to the problem.
502 You can see that refactoring V2 introduces overly specialized functions like `dijkstra_special()` or
503 `compute_min_moves_opposite_parity()`. In comparison, refactoring V1 generates only general
504 versions of these functions (e.g. `dijkstra()`).

```

1
2 #==== NEW HELPER FUNCTIONS ====
3 def read_ints():
4     return list(map(int, input().split()))
5
6 def build_adj_undirected(n, edges):
7     adj = [[] for _ in range(n)]
8     for u, v, w in edges:
9         adj[u].append((v, w))
10        adj[v].append((u, w))
11    return adj
12
13 def dijkstra(adj, src):
14     from heapq import heappush, heappop
15     INF = 10**18
16     n = len(adj)
17     dist = [INF]*n
18     parent = [-1]*n
19     dist[src] = 0
20     heap = [(0, src)]
21     while heap:
22         d, u = heappop(heap)
23         if d > dist[u]:
24             continue
25         for v, w in adj[u]:
26             nd = d + w
27             if nd < dist[v]:
28                 dist[v] = nd
29                 parent[v] = u
30                 heappush(heap, (nd, v))
31    return dist, parent
32
33 def reconstruct_path(parent, dest):
34     path = []
35     u = dest
36     while u != -1:
37         path.append(u+1)
38         u = parent[u]
39    return path[::-1]
40
41 def multi_source_bfs(neighbors, sources):
42     from collections import deque
43     n = len(neighbors)
44     dist = [-1]*n
45     dq = deque()
46     for u in sources:
47         if dist[u] == -1:
48             dist[u] = 0
49             dq.append(u)
50     while dq:
51         u = dq.popleft()
52         for v in neighbors[u]:
53             if dist[v] == -1:
54                 dist[v] = dist[u] + 1
55                 dq.append(v)
56    return dist

```

Listing 10: Version 1, New Helpers

```

1  ##### PROGRAM: node_16:cc_python_16 #####
2
3  from codebank import *
4
5  def main():
6      import heapq
7      n, m = read_ints()
8      edges = [(u-1, v-1, w) for u, v, w in (read_ints() for _ in range(m))]
9      adj = build_adj_undirected(n, edges)
10     INF = 10**20
11     dist = [INF]*n
12     dist[0] = 0
13     last_w = [0]*n
14     heap = [(0, 0)]
15     while heap:
16         d, u = heapq.heappop(heap)
17         if d > dist[u]:
18             continue
19         # record last edges
20         for v, w in adj[u]:
21             last_w[v] = w
22         # expand two-edge moves
23         for v, w1 in adj[u]:
24             tw = last_w[v]
25             for x, w2 in adj[v]:
26                 nd = d + (tw + w2)**2
27                 if nd < dist[x]:
28                     dist[x] = nd
29                     heapq.heappush(heap, (nd, x))
30
31     out = []
32     for x in dist:
33         out.append(str(x if x < INF else -1))
34     print(" ".join(out))
35
36 if __name__ == "__main__":
37     main()
38
39 ##### PROGRAM: node_17:cc_python_17 #####
40
41 from codebank import *
42
43 def main():
44     n, m = read_ints()
45     edges = [(u-1, v-1, w) for u, v, w in (read_ints() for _ in range(m))]
46     adj = build_adj_undirected(n, edges)
47     dist, parent = dijkstra(adj, 0)
48     if dist[n-1] >= 10**18:
49         print(-1)
50     else:
51         path = reconstruct_path(parent, n-1)
52         print(*path)
53
54 if __name__ == "__main__":
55     main()
56
57 ##### PROGRAM: node_19:cc_python_19 #####
58
59 from codebank import *
60
61 def main():
62     n = int(input())
63     a = read_ints()
64     # build reversed graph: for each move i->j, add edge j->i
65     neighbors = [[] for _ in range(n)]
66     for i, val in enumerate(a):
67         for j in (i - val, i + val):
68             if 0 <= j < n:
69                 neighbors[j].append(i)
70     # BFS from all even and all odd positions separately
71     even_sources = [i for i, val in enumerate(a) if val % 2 == 0]
72     odd_sources = [i for i, val in enumerate(a) if val % 2 == 1]
73     dist_even = multi_source_bfs(neighbors, even_sources)
74     dist_odd = multi_source_bfs(neighbors, odd_sources)
75     # for odd a[i], answer is dist to nearest even => dist_even; else dist_odd
76     ans = [dist_even[i] if a[i] % 2 == 1 else dist_odd[i] for i in range(n)]
77     print(*ans)
78
79 if __name__ == "__main__":
80     main()

```

Listing 11: Version 1, Refactored Programs

```

1 #
2 #==== NEW HELPER FUNCTIONS ====
3 def read_ints():
4     return list(map(int, input().split()))
5
6 def build_undirected_weighted_graph(n, m):
7     from collections import defaultdict
8     adj = defaultdict(list)
9     for _ in range(m):
10         u, v, w = read_ints()
11         u -= 1; v -= 1
12         adj[u].append((v, w))
13         adj[v].append((u, w))
14     return adj
15
16 def dijkstra(adj, src, n):
17     import heapq
18     INF = 10**18
19     dist = [INF]*n
20     parent = [-1]*n
21     visited = [False]*n
22     dist[src] = 0
23     heap = [(0, src)]
24     while heap:
25         d, u = heapq.heappop(heap)
26         if visited[u]:
27             continue
28         visited[u] = True
29         for v, w in adj.get(u, ()):
30             nd = d + w
31             if nd < dist[v]:
32                 dist[v] = nd
33                 parent[v] = u
34             heapq.heappush(heap, (nd, v))
35     return dist, parent
36
37 def reconstruct_path(parent, dest):
38     path = []
39     while dest != -1:
40         path.append(dest+1)
41         dest = parent[dest]
42     return path[::-1]
43
44 def dijkstra_special(e, n, src):
45     import heapq
46     INF = 10**18
47     d = [INF]*n
48     d[src] = 0
49     heap = [(0, src)]
50     while heap:
51         cd, v = heapq.heappop(heap)
52         if cd > d[v]:
53             continue
54         td = {}
55         for u, w in e.get(v, ()):
56             td[u] = w
57         for u, w1 in td.items():
58             for x, w2 in e.get(u, ()):
59                 cost = cd + (w1 + w2)**2
60                 if cost < d[x]:
61                     d[x] = cost
62                     heapq.heappush(heap, (cost, x))
63     return d
64
65 def compute_min_moves_opposite_parity(a):
66     from collections import deque
67     n = len(a)
68     go = [[] for _ in range(n)]
69     ans = [-1]*n
70     q = deque()
71     for i, val in enumerate(a):
72         for j in (i - val, i + val):
73             if 0 <= j < n:
74                 if (a[j] % 2) != (val % 2):
75                     ans[i] = 1
76                     q.append(i)
77                     break
78             else:
79                 go[j].append(i)
80     while q:
81         u = q.popleft()
82         for v in go[u]:
83             if ans[v] == -1:
84                 ans[v] = ans[u] + 1
85                 q.append(v)
86     return ans

```

Listing 12: Version 2, New Helpers

```

1 #
2 # ##### PROGRAM: node_16:cc_python_16 #####
3
4 from codebank import *
5
6 def main():
7     n, m = read_ints()
8     e = {}
9     for _ in range(m):
10         u, v, w = read_ints()
11         u -= 1; v -= 1
12         e.setdefault(u, []).append((v, w))
13         e.setdefault(v, []).append((u, w))
14     d = dijkstra_special(e, n, 0)
15     print(" ".join(str(-1 if x >= 10**18 else int(x)) for x in d))
16
17 if __name__ == "__main__":
18     main()
19
20 # ##### PROGRAM: node_17:cc_python_17 #####
21
22 from codebank import *
23
24 def main():
25     n, m = read_ints()
26     adj = build_undirected_weighted_graph(n, m)
27     dist, parent = dijkstra(adj, 0, n)
28     if dist[n-1] >= 10**18:
29         print(-1)
30     else:
31         path = reconstruct_path(parent, n-1)
32         print(" ".join(map(str, path)))
33
34 if __name__ == "__main__":
35     main()
36
37 # ##### PROGRAM: node_19:cc_python_19 #####
38
39 from codebank import *
40
41 def main():
42     n = int(input())
43     a = read_ints()
44     ans = compute_min_moves_opposite_parity(a)
45     print(" ".join(map(str, ans)))
46
47 if __name__ == "__main__":
48     main()

```

Listing 13: Version 2, Refactored Programs

```

1 from datapipe.core import (
2     tumbling_window,
3     sliding_window as _sliding_window,
4     add_serializer,
5     throttle_upstream as _throttle_upstream,
6     watermark_event_time as _watermark_event_time,
7     ...
8 )
9 ...
10 ...
11
12 def throttle_upstream(max_size):
13     """
14     Apply backpressure to slow data ingestion if downstream stages are overloaded.
15
16     Args:
17         max_size: maximum queue size or rate limit
18
19     Returns:
20         Decorator function
21     """
22     def decorator(func):
23         from functools import wraps
24
25         @wraps(func)
26         def queue_wrapper(q, *args, **kwargs):
27             try:
28                 size = q.qsize()
29                 if size > max_size:
30                     import time
31                     time.sleep(0.01)
32             except Exception:
33                 pass
34             return func(q, *args, **kwargs)
35
36         return queue_wrapper
37
38     return decorator
39
40 def watermark_event_time(events, allowed_lateness):
41     """
42     Assign event-time watermarks to handle late data correctly.
43
44     Args:
45         events: list of dicts with timestamp
46         allowed_lateness: seconds of allowed lateness
47
48     Returns:
49         Events with watermark annotations
50     """
51     # Ensure we return the appropriate format with is_late field
52     result = []
53     for e in events:
54         tagged = dict(e)
55         max_ts = max(ev['timestamp'] for ev in events)
56         watermark = max_ts - allowed_lateness
57         tagged['watermark'] = watermark
58         tagged['is_late'] = e['timestamp'] < watermark
59         result.append(tagged)
60     return result
61 ...

```

Listing 14: Claude fails to use imports and instead re-implements the function.