

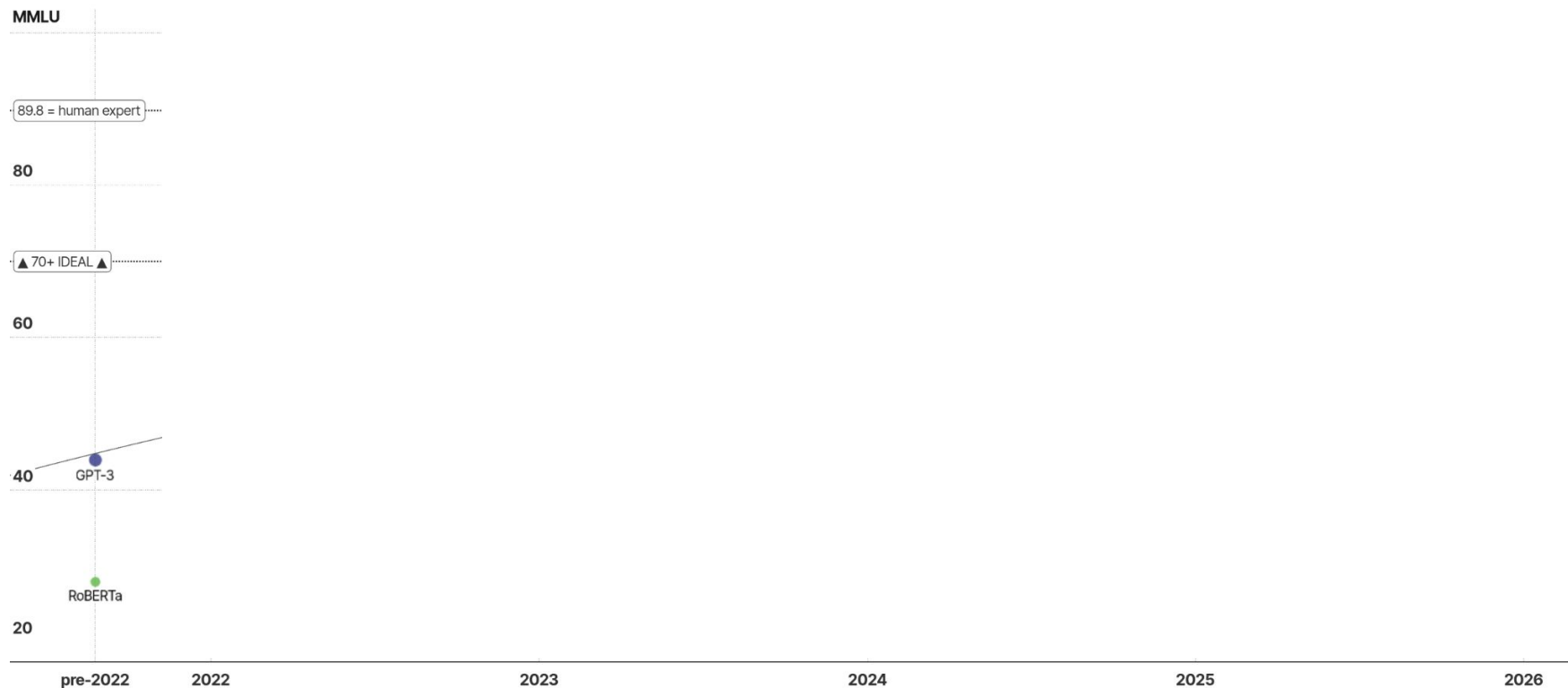
A Catalog of Data Smells for Coding Tasks

Antonio Vitale, Rocco Oliveto, Simone Scalabrino

ICSE Journal First
Transactions on Software Engineering and Methodology (TOSEM)



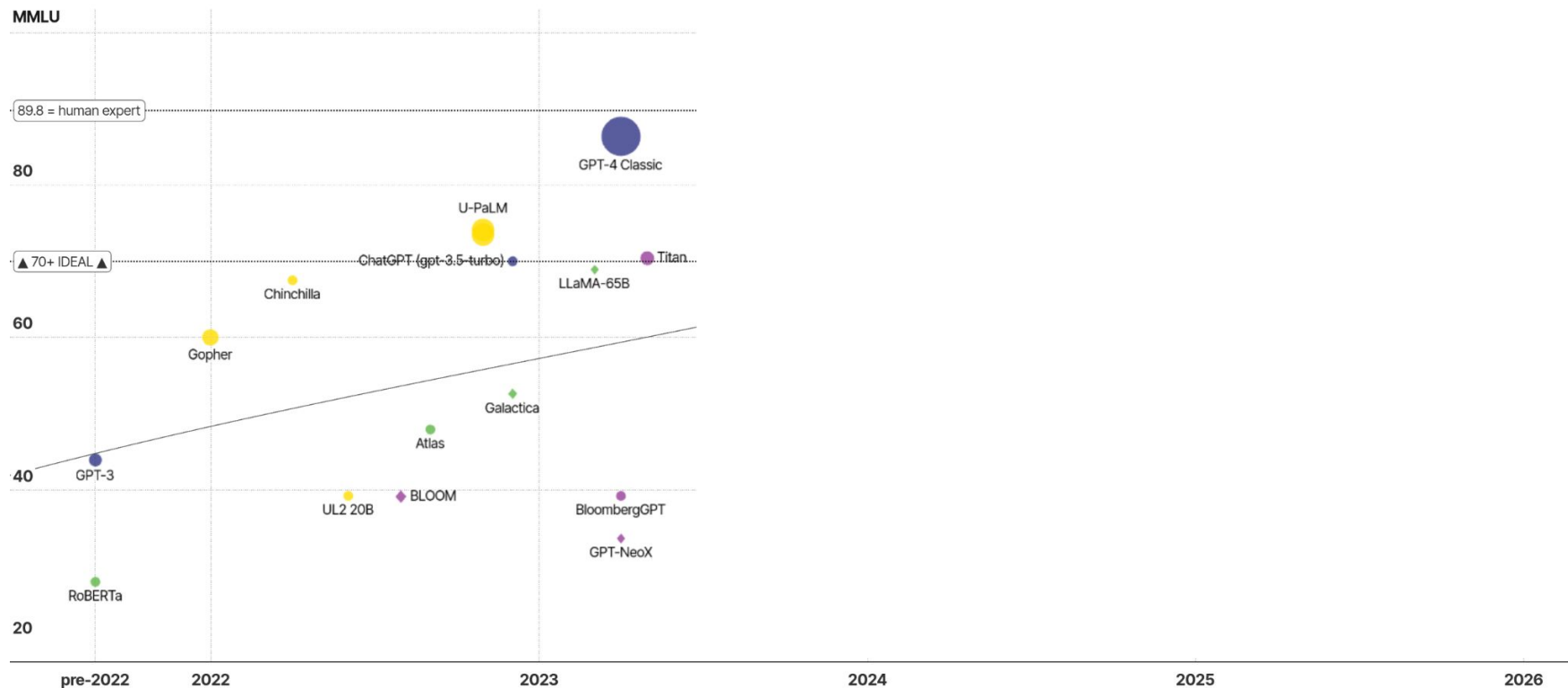
BACKGROUND



David McCandless, Tom Evans, Paul Barton
Informationisbeautiful // Jan 2026

MMLU = benchmark for measuring LLM capabilities
* = parameters undisclosed // source: [LifeArchitect](#) // [data](#)

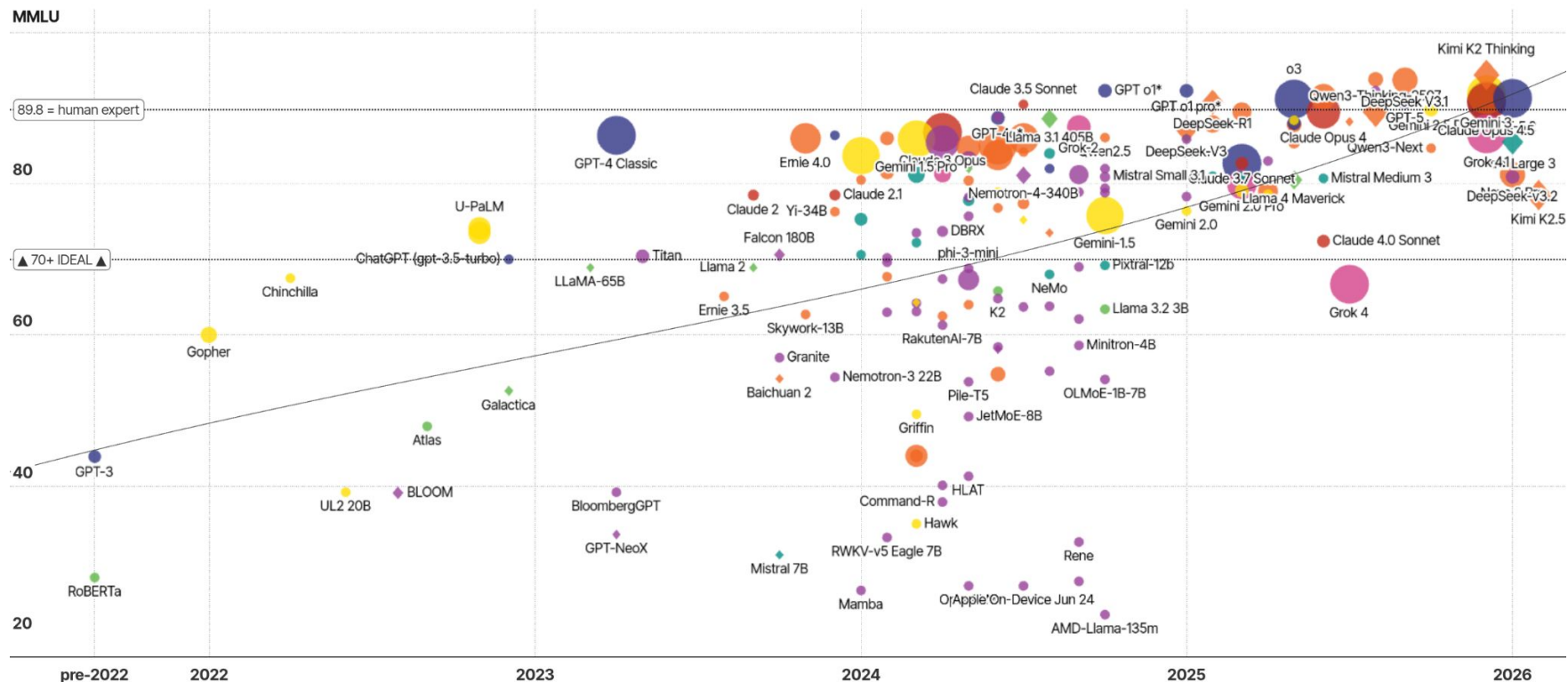
BACKGROUND



David McCandless, Tom Evans, Paul Barton
Informationisbeautiful // Jan 2026

MMLU = benchmark for measuring LLM capabilities
* = parameters undisclosed // source: [LifeArchitect](#) // [data](#)

BACKGROUND



BACKGROUND

arXiv:2002.08155v4 [cs.CL] 18 Sep 2020

CodeBERT:
A Pre-Trained Model for Programming and Natural Languages
Zhengyu Feng¹, Daya Guo², Deza Tang¹, Nan Duan¹, Xuecheng Feng¹,
Ming Gong¹, Lijian Shao¹, Bing Qin¹, Ting Lai¹, Daxin Jiang¹, Ming Zhou¹
¹ Research Center for Social Computing and Information Retrieval, Harbin Institute of Technology, China
² The School of Data and Computer Science, Sun Yat-sen University, China
³ Microsoft Research Technology Center Asia, Beijing, China
[fzyfeng, mfgong, qibing, lili@jitai.edu.cn, mzhou]
[guoddy2004112, tyysu.edu.cn]
[dutanang, nanduan, nfguo, lishao, djdang, xingzhou@microsoft.com]

Abstract

We present CodeBERT, a bilingual pre-trained model for programming language (PL) and natural language (NL). CodeBERT learns general-purpose representations that support downstream NLP applications such as natural language code search, code documentation generation, etc. We design CodeBERT with Transformer-based neural architecture, and train it with a hybrid objective function that incorporates the pre-training task of masked token detection, which is to detect plausible alternatives sampled from generators. The model is so called “bilingual”. Data of NL, PL pairs and “natural” data, where the former provides input tokens for model training while the latter helps to learn better generators. We evaluate CodeBERT on two NLP applications by fine-tuning model parameters. Results show that CodeBERT achieves state-of-the-art performance on both natural language code search and code documentation generation. Furthermore, we investigate what type of knowledge is learned in CodeBERT, we construct a dataset for NL, PL, pairing, and release it as a pre-training data source. Evaluation of pre-trained models and final results show that CodeBERT performs better than previous pre-trained models on NL, PL pairs.¹

1 Introduction

Large pre-trained models such as ELMo (Peters et al., 2018), GPT (Radford et al., 2018), BERT (Devlin et al., 2019), XLNet (Yang et al., 2019)

¹“Data Release” while this author was at Microsoft Research Asia.
²“All the codes and data are available at <https://github.com/microsoft/codebert>”

and RoBERTa (Liu et al., 2019) have dramatically improved the state-of-the-art on a variety of natural language processing (NLP) tasks. These pre-trained models have effectively extracted representations from massive unlabeled text optimized by self-supervised objectives, such as masked language modeling, which predicts the original masked word from an artificially masked input sequence. The success of pre-trained models in NLP also opens a range of rich model pre-training models, such as ViLBERT (Liu et al., 2019) for language-image and VideoBERT (Gao et al., 2019) for language-video, which are learned from bi-modal data sets or language-image pairs with bi-modal self-supervised objectives.

In this work, we present CodeBERT, a bilingual pre-trained model for natural language (NL) and programming language (PL) like Python, Java, JavaScript, etc. CodeBERT captures the semantic connection between natural language and programming language, and encodes general-purpose representations that can broadly support NLP, understanding tasks (e.g., natural language code search) and generation tasks (e.g., code documentation generation). It is developed with the multi-layer Transformer (Vaswani et al., 2017), which is adapted for a majority of large pre-trained models. In order to make use of both bilingual instances of NL, PL pairs and large amount of available natural codes, we train CodeBERT with a hybrid objective function, including masked natural language modeling (Devlin et al., 2019) and masked token detection (Clark et al., 2020), where natural codes help to learn better generators for producing better alternative tokens for the later objective.

We train CodeBERT from Github code repository.

2020

BACKGROUND

arXiv:2002.08155v4 [cs.CL] 18 Sep 2020

CodeBERT: A Pre-Trained Model for Programming and Natural Languages

Zhuncheng Feng¹, Daya Guo², Dehao Tang¹, Nan Du¹, Xuecheng Feng¹, Ming Zhou¹, Jindan Zhou¹, Bing Qin¹, Ting Liu¹, Daxin Jiang¹, Ming Zhou¹
¹Research Center for Social Computing and Information Retrieval, Harbin Institute of Technology, China
²Microsoft Research Asia, Beijing, China
³Microsoft Research Technology Center Asia, Beijing, China
[fzfeng, dguo, dtang, jzhang, mingzhou, jzhang, mingzhou]@msrc.hit.edu.cn
[guo, tang, nan, du, guo, jzhang, mingzhou]@microsoft.com

Abstract

We present CodeBERT, a bilingual pre-trained model for programming language (PL) and natural language (NL). CodeBERT learns general-purpose representations that support downstream NLP applications such as natural language code search, code documentation generation, etc. We design CodeBERT with Transformer-based neural networks, and train it with a hybrid objective function that incorporates the pre-training task of masked token detection, which is to predict possible alternatives completed from fragments. The model is so-called “bilingual” because of its PL and NL parts and “universal” data, where the former provides input tokens for model pre-training. We evaluate CodeBERT on two NLP applications by fine-tuning model parameters. Results show that CodeBERT achieves state-of-the-art performance on both natural language code search and code documentation generation. Furthermore, we investigate how type of knowledge is learned in CodeBERT, we conduct a dataset for NL, PL, parsing, and related tasks to compare the quality of representations of pre-trained models and find that CodeBERT performs better than the previous pre-trained models on NL.

1 Introduction

Large pre-trained models such as BERT (Bert et al., 2019), GPT (Radford et al., 2018), PEG (Devlin et al., 2018), XLNet (Yang et al., 2019), etc. have shown that this model can be used as a backbone for many tasks and data are available.

We train CodeBERT from Github code repos.

IntelliCode Composer: Code Generation using Transformer

Alexey Svyatkovskiy¹
Robyn Wang²
Alexey Svyatkovskiy¹
Shengyu Fu¹
Microsoft
Redmond, WA, USA
alexsv@microsoft.com

Shao Kun Deng¹
Maoen Wu¹
Jiahua Sun¹
Neel Sundaresan¹
Microsoft
Redmond, WA, USA
neelsun@microsoft.com

ABSTRACT

In software development through integrated development environments (IDEs), code completion is one of the most widely used features. Nevertheless, majority of popular code completion solutions only support completion of methods and APIs or sequences. In this paper, we introduce IntelliCode Composer, a general-purpose intelligent code completion tool which is capable of predicting sequence of code tokens of arbitrary types, generating up to entire lines of externally created code. It leverages state-of-the-art generative transformer model trained on a collection of source code as PL tokens. It integrates and leverages pre-training language. IntelliCode Composer is implemented as a third-party extension to Visual Studio Code IDE and Atom IDE.

KEYWORDS

Software and its engineering • Integrated and cloud-based software environments • Computing methodologies • Natural language processing

Code completion, natural languages, introduction of software

ACM Reference Format:

Alexey Svyatkovskiy, Robyn Wang, Shengyu Fu, and Neel Sundaresan.

IntelliCode Composer: Code Generation using Transformer. In *Proceedings of the ACM Conference on Computer Science Research*, 2020.

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

https://doi.org/10.1145/3388888.3388888

A Transformer-based Approach for Source Code Summarization

Wael Uddin Ahmad¹
University of California, Los Angeles
waelu@ucla.edu

Balabkhi Ray¹
Columbia University
rayb18@cs.columbia.edu

Sakshat Chakrabarti²
Columbia University
sakshat@cs.columbia.edu

Kai Wei Chung¹
University of California, Los Angeles
kchung@ucla.edu

Abstract

Generating a readable summary that describes the functionality of a program is known as source code summarization. In this task, we propose a novel approach to generate a summary of source code by leveraging code representation by modeling the sequence-to-sequence relationship between code tokens in a sequence-to-sequence framework. To learn code representation for summarization, we propose a Transformer model that uses a self-attention mechanism and has shown to be effective in capturing long-range dependencies. In this work, we show that despite the approach is simple, it outperforms the state-of-the-art approaches by a significant margin. We perform extensive qualitative analysis that reveal several important findings, e.g., the effective modeling of source code tokens, precise analysis, while summarization significantly improves the summarization performance. We have our code publicly available to facilitate future research.

1 Introduction

Program comprehension is an indispensable ingredient of software development and maintenance (Chen et al., 2018). A natural language summary of source code facilitates program comprehension by reducing developer efforts significantly (Scherer et al., 2010). Source code summarization is a task that aims to generate a summary that describes the functionality of a program.

With the advancement of deep learning and availability of large-scale data, there has been a vast number of open-source representations, automatic code summarization tools, and summarization tools from researchers. Most of the state-of-the-art approaches are based on the sequence-to-sequence framework. One of the initial works, Pyrit (Dallmeier et al., 2016), used a sequence-to-sequence framework to generate a summary of source code.

Many efforts to represent the individual code tokens and combine them with a Recurrent Neural Network (RNN) as an attention mechanism to generate a natural language summary. Subsequent works, Wang and Zhou (2018), He et al. (2018a), adopted the traditional RNN-based sequence-to-sequence model (Scherer et al., 2010) with an explicit first-order dependency to generate a summary. To learn code representation for summarization, we propose a Transformer model that uses a self-attention mechanism and has shown to be effective in capturing long-range dependencies. In this work, we show that despite the approach is simple, it outperforms the state-of-the-art approaches by a significant margin. We perform extensive qualitative analysis that reveal several important findings, e.g., the effective modeling of source code tokens, precise analysis, while summarization significantly improves the summarization performance. We have our code publicly available to facilitate future research.

arXiv:2005.00653v1 [cs.SE] 1 May 2020

2020

2020

2020

arXiv:2002.08155v4 [cs.CL] 18 Sep 2020

Abstract

pre-trained CoBERT, a bi-directional encoder, is used for preprocessing language (NL) and image (I) inputs. The pre-trained model provides a word-level representation that captures the semantic information of the natural language side and document-level information on the image side. CoBERT with frequency-based neural architecture search (NAS) is used to find the best model for the task of predicting the missing link of replaced tokens in images, which is direct evidence of the model's ability to understand the meaning behind the "visual" data. The models are trained both "blindly" data-wise and "wisely" with the help of the NAS framework. The former provides rich labels for model training, while the latter helps to find the best model for the task of image-to-text generation. We evaluate CoBERT on the task of image-to-text generation with different parameters. Results show that CoBERT can be used to generate natural language from images and vice versa. The model also demonstrates generalization. In particular, CoBERT can generate NL from I, even though it was not trained for this task. CoBERT is our current best model for image-to-text generation, as shown by the parameters of pre-trained models and the results of the NAS framework. The model is better than previous pre-trained models on NL.

¹ All the codes and data are available at <https://github.com/microsoft/code42>

We train CodeBERT from Github code repository-

Shao Kun Deng*
Microsoft
Redmond, WA, USA
shudo@microsoft.com

Neel Sundaresan
Microsoft
Redmond, WA, USA
neels@microsoft.com

Xiv:2005.08025v2 [cs.CL] 29 Oct 2020

ABSTRACT In software development through integrated development environments (IDEs), code completion is one of the most widely used features. Nevertheless, majority of integrated development environments only support completion of methods and APIs, or generate suggestions for code completion. In this paper, we propose a novel code completion tool that can complete any code by predicting sequences of code tokens of arbitrary types, generating up to *m* lines of syntactically correct code. It leverages state of the art generative transformer model trained on 1.2 billion lines of source code in Python, C#, JavaScript and TypeScript programming languages. IntelliCode is deployed as a cloud-based service that can be used by any IDE. It can be used to generate parallel implementation of the beam search decoder, and structural graph optimizations to meet real-time completion suggestions in queries in the Visual Studio Code IDE and Atom text editor.

Our best model yields an average edit similarity of 86.7% and perplexity of 1.82 for Python programming language.

KEYWORDS

KEYWORDS
Code completion, neural networks, naturalness of software

ACM Reference Format:
Alamy Sayakulchit, Shao Xia Dong, Shengyu Fu, and Naveel Sundaresan. 2020. IntelliCode Composer: Code Generation using Transformers. In *Proceedings of the 20th ACM Joint European Software Engineering Conference and Symposium on the Foundations of Software Engineering (ESEC/FSE '20)*, November 8–13, 2020, Virtual Event, USA. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3606891.3710758>

Permission to make digital or hard copies of all or part of this work for personal or internal use, or the personal or internal use of specific clients, is granted by ACM for users registered with ACM, provided that the fee of \$12.00 per copy is paid directly to ACM. For those organizations that have been granted a corporate rate of payment, a separate system of payment has been arranged. The fee code for users registered with ACM is 0001-0706/2010/04 \$12.00. ACM is not authorized to make copies for general distribution, for advertising or promotional purposes, for creating new collective works, or for resale.

[illegible]

INTRODUCTION

Machine learning has shown a great promise towards improving automated software engineering across all stages. Some of the early applications of machine learning of source code include code classification [1], code summarization [2], code completion [3–5] and code generation [6].

There are numerous code completion systems capable of offering suggestions to the user as he/she types. However, not all of them are the correct answer [12–14]. Majority of empirical completion systems would, however, only work when the name of the method or API call is already typed in, thus leaving the task of completing the rest of the code to the user.

In this paper, we introduce *IntelliCode* (a general purpose code completion framework, capable of generating code snippets of arbitrary lines/types, including local variables, method calls, imports, package names, etc.) that can be used by IDEs and developers. *IntelliCode* serves as a universal programmatic language modeling tool, effectively generating syntactically correct code snippets for a given context. *IntelliCode* is implemented as an entire line of code instead of a snippet, thus saving the user from the need of manually constructing a snippet with the proper experience inspired from Visual Studio Code [15]. The user can

The nature of the problem of code sequence completion makes statistical language modeling appear a promising starting point. To predict a whole line of source code tokens given an existing code sequence, we need to model the probability of a sequence of tokens $\{x_1, \dots, x_n\} \in V^n$, V is the vocabulary, conditioned on a sequence of previously issued tokens $\{x_1, \dots, x_{n-1}\} \in V^{n-1}$ of code snippet C .

The main contributions of the paper are as follows: (i) we introduce and present a neural layer generative transformer model for code sequence completion; (ii) we compare the performance of the proposed model on a large number of synthetic and real-world code completion tasks (i.e. sections 4 and 5), comparing it to the GPT-2 (OpenAI, 2019) and the state-of-the-art code completion models; (iii) we propose and display a novel end to end code sequence completion system and build a full-fledged code completion system on the GPT-2 and an additional code completion model (see section 7.1 and 7.2); (iv) we evaluate the quality of language model pretraining of GPT-2 using perplexity, showing that our best model achieves a perplexity of 1.82; we also show that the *triforce* Code Composer achieves an average edit distance

Code Compose achieves an average edit similitude

Saikat Chakraborty
Columbia University
saikatc@cs.columbia.edu

Kai-Wei Chang

Generating a readable summary that describes the functionality of a program is known as *program summarization*. This is a challenging task because of the complexity of the code. Existing work on program summarization has been done by modeling the relationships between code elements in a program and their corresponding summaries. To learn code representation for summarization, existing work has used neural networks, but used a self-attention mechanism and a hierarchical structure to capture the dependencies in the code. In this work, we show that deep learning can be used to generate a summary of the code by using a recurrent neural network (RNN) based sequence model and combining this with a Recurrent Neural Network (RNN) using an attention mechanism to generate a summary of the code. We evaluate our model on a dataset of code snippets and their corresponding summaries (Liang and Zou, 2018; Bi et al., 2018) and compare our model with existing work on program summarization (Zhang and Chen, 2014; Bi et al., 2018) using an attention mechanism (Liang et al., 2015) on two different datasets.

The RNN-based sequence models have different limitations in learning source code representations. First, the RNN-based models are limited in the length of the sequence of code snippets. Second, some code snippets are very long and contain many dependencies. Third, the RNN-based models are limited in the length of the sequence of code snippets. Fourth, the RNN-based models are limited in the length of the sequence of code snippets. Finally, the RNN-based models are limited in the length of the sequence of code snippets.

Introduction

Program comprehension is an indispensable ingredient of software development and maintenance (Kao et al., 2004). A natural language summary (NLS) is a textual representation of a program that aids in reducing developers' efforts significantly (Sridharan et al., 1999). Source code summarization is a natural language processing (NLP) problem that describes the functionality of a program.

Due to the advancement of deep learning and the availability of large-scale data through a vast number of open-source repositories, natural language processing has become a hot research area from researchers. Most of the natural approaches generate source code summaries in a sentence or a few lines (Sridharan et al., 1999; Wang and Yorks et al., 2010) trained on code snippets.

Program comprehension is an indispensable ingredient of software development and maintenance (Kao et al., 2004). A natural language summary (NLS) is a textual representation of a program that aids in reducing developers' efforts significantly (Sridharan et al., 1999). Source code summarization is a natural language processing (NLP) problem that describes the functionality of a program.

Due to the advancement of deep learning and the availability of large-scale data through a vast number of open-source repositories, natural language processing has become a hot research area from researchers. Most of the natural approaches generate source code summaries in a sentence or a few lines (Sridharan et al., 1999; Wang and Yorks et al., 2010) trained on code snippets.

Program comprehension is an indispensable ingredient of software development and maintenance (Wang et al. 2016). The development of a new program (or source code) facilitates program comprehension by reducing developers' effort significantly (Wang et al. 2016). However, the development of a new program does not always reduce the task of creating maintainable summaries that describe the functionality of a program (Wang et al. 2016).

With the advancement of deep learning and the availability of large-scale data, deep learning has been widely used in program comprehension (Wang et al. 2016). Deep learning has drawn attention from the research community because it can process *generic* source code representations in a sequence-to-sequence fashion. One of the initial works in this area (Chen et al. 2016) embedded source code as machine translation (Wang et al. 2016; Wang et al. 2017; Wang et al. 2018; Wang et al. 2019; Wang et al. 2018c; Wang et al. 2018b; Wang et al. 2018a; Wang et al. 2018d; Wang et al. 2018e; Wang et al. 2018f; Wang et al. 2018g; Wang et al. 2018h; Wang et al. 2018i; Wang et al. 2018j; Wang et al. 2018k; Wang et al. 2018l; Wang et al. 2018m; Wang et al. 2018n; Wang et al. 2018o; Wang et al. 2018p; Wang et al. 2018q; Wang et al. 2018r; Wang et al. 2018s; Wang et al. 2018t; Wang et al. 2018u; Wang et al. 2018v; Wang et al. 2018w; Wang et al. 2018x; Wang et al. 2018y; Wang et al. 2018z; Wang et al. 2019a; Wang et al. 2019b; Wang et al. 2019c; Wang et al. 2019d; Wang et al. 2019e; Wang et al. 2019f; Wang et al. 2019g; Wang et al. 2019h; Wang et al. 2019i; Wang et al. 2019j; Wang et al. 2019k; Wang et al. 2019l; Wang et al. 2019m; Wang et al. 2019n; Wang et al. 2019o; Wang et al. 2019p; Wang et al. 2019q; Wang et al. 2019r; Wang et al. 2019s; Wang et al. 2019t; Wang et al. 2019u; Wang et al. 2019v; Wang et al. 2019w; Wang et al. 2019x; Wang et al. 2019y; Wang et al. 2019z; Wang et al. 2020a; Wang et al. 2020b; Wang et al. 2020c; Wang et al. 2020d; Wang et al. 2020e; Wang et al. 2020f; Wang et al. 2020g; Wang et al. 2020h; Wang et al. 2020i; Wang et al. 2020j; Wang et al. 2020k; Wang et al. 2020l; Wang et al. 2020m; Wang et al. 2020n; Wang et al. 2020o; Wang et al. 2020p; Wang et al. 2020q; Wang et al. 2020r; Wang et al. 2020s; Wang et al. 2020t; Wang et al. 2020u; Wang et al. 2020v; Wang et al. 2020w; Wang et al. 2020x; Wang et al. 2020y; Wang et al. 2020z; Wang et al. 2021a; Wang et al. 2021b; Wang et al. 2021c; Wang et al. 2021d; Wang et al. 2021e; Wang et al. 2021f; Wang et al. 2021g; Wang et al. 2021h; Wang et al. 2021i; Wang et al. 2021j; Wang et al. 2021k; Wang et al. 2021l; Wang et al. 2021m; Wang et al. 2021n; Wang et al. 2021o; Wang et al. 2021p; Wang et al. 2021q; Wang et al. 2021r; Wang et al. 2021s; Wang et al. 2021t; Wang et al. 2021u; Wang et al. 2021v; Wang et al. 2021w; Wang et al. 2021x; Wang et al. 2021y; Wang et al. 2021z; Wang et al. 2022a; Wang et al. 2022b; Wang et al. 2022c; Wang et al. 2022d; Wang et al. 2022e; Wang et al. 2022f; Wang et al. 2022g; Wang et al. 2022h; Wang et al. 2022i; Wang et al. 2022j; Wang et al. 2022k; Wang et al. 2022l; Wang et al. 2022m; Wang et al. 2022n; Wang et al. 2022o; Wang et al. 2022p; Wang et al. 2022q; Wang et al. 2022r; Wang et al. 2022s; Wang et al. 2022t; Wang et al. 2022u; Wang et al. 2022v; Wang et al. 2022w; Wang et al. 2022x; Wang et al. 2022y; Wang et al. 2022z; Wang et al. 2023a; Wang et al. 2023b; Wang et al. 2023c; Wang et al. 2023d; Wang et al. 2023e; Wang et al. 2023f; Wang et al. 2023g; Wang et al. 2023h; Wang et al. 2023i; Wang et al. 2023j; Wang et al. 2023k; Wang et al. 2023l; Wang et al. 2023m; Wang et al. 2023n; Wang et al. 2023o; Wang et al. 2023p; Wang et al. 2023q; Wang et al. 2023r; Wang et al. 2023s; Wang et al. 2023t; Wang et al. 2023u; Wang et al. 2023v; Wang et al. 2023w; Wang et al. 2023x; Wang et al. 2023y; Wang et al. 2023z; Wang et al. 2024a; Wang et al. 2024b; Wang et al. 2024c; Wang et al. 2024d; Wang et al. 2024e; Wang et al. 2024f; Wang et al. 2024g; Wang et al. 2024h; Wang et al. 2024i; Wang et al. 2024j; Wang et al. 2024k; Wang et al. 2024l; Wang et al. 2024m; Wang et al. 2024n; Wang et al. 2024o; Wang et al. 2024p; Wang et al. 2024q; Wang et al. 2024r; Wang et al. 2024s; Wang et al. 2024t; Wang et al. 2024u; Wang et al. 2024v; Wang et al. 2024w; Wang et al. 2024x; Wang et al. 2024y; Wang et al. 2024z; Wang et al. 2025a; Wang et al. 2025b; Wang et al. 2025c; Wang et al. 2025d; Wang et al. 2025e; Wang et al. 2025f; Wang et al. 2025g; Wang et al. 2025h; Wang et al. 2025i; Wang et al. 2025j; Wang et al. 2025k; Wang et al. 2025l; Wang et al. 2025m; Wang et al. 2025n; Wang et al. 2025o; Wang et al. 2025p; Wang et al. 2025q; Wang et al. 2025r; Wang et al. 2025s; Wang et al. 2025t; Wang et al. 2025u; Wang et al. 2025v; Wang et al. 2025w; Wang et al. 2025x; Wang et al. 2025y; Wang et al. 2025z; Wang et al. 2026a; Wang et al. 2026b; Wang et al. 2026c; Wang et al. 2026d; Wang et al. 2026e; Wang et al. 2026f; Wang et al. 2026g; Wang et al. 2026h; Wang et al. 2026i; Wang et al. 2026j; Wang et al. 2026k; Wang et al. 2026l; Wang et al. 2026m; Wang et al. 2026n; Wang et al. 2026o; Wang et al. 2026p; Wang et al. 2026q; Wang et al. 2026r; Wang et al. 2026s; Wang et al. 2026t; Wang et al. 2026u; Wang et al. 2026v; Wang et al. 2026w; Wang et al. 2026x; Wang et al. 2026y; Wang et al. 2026z; Wang et al. 2027a; Wang et al. 2027b; Wang et al. 2027c; Wang et al. 2027d; Wang et al. 2027e; Wang et al. 2027f; Wang et al. 2027g; Wang et al. 2027h; Wang et al. 2027i; Wang et al. 2027j; Wang et al. 2027k; Wang et al. 2027l; Wang et al. 2027m; Wang et al. 2027n; Wang et al. 2027o; Wang et al. 2027p; Wang et al. 2027q; Wang et al. 2027r; Wang et al. 2027s; Wang et al. 2027t; Wang et al. 2027u; Wang et al. 2027v; Wang et al. 2027w; Wang et al. 2027x; Wang et al. 2027y; Wang et al. 2027z; Wang et al. 2028a; Wang et al. 2028b; Wang et al. 2028c; Wang et al. 2028d; Wang et al. 2028e; Wang et al. 2028f; Wang et al. 2028g; Wang et al. 2028h; Wang et al. 2028i; Wang et al. 2028j; Wang et al. 2028k; Wang et al. 2028l; Wang et al. 2028m; Wang et al. 2028n; Wang et al. 2028o; Wang et al. 2028p; Wang et al. 2028q; Wang et al. 2028r; Wang et al. 2028s; Wang et al. 2028t; Wang et al. 2028u; Wang et al. 2028v; Wang et al. 2028w; Wang et al. 2028x; Wang et al. 2028y; Wang et al. 2028z; Wang et al. 2029a; Wang et al. 2029b; Wang et al. 2029c; Wang et al. 2029d; Wang et al. 2029e; Wang et al. 2029f; Wang et al. 2029g; Wang et al. 2029h; Wang et al. 2029i; Wang et al. 2029j; Wang et al. 2029k; Wang et al. 2029l; Wang et al. 2029m; Wang et al. 2029n; Wang et al. 2029o; Wang et al. 2029p; Wang et al. 2029q; Wang et al. 2029r; Wang et al. 2029s; Wang et al. 2029t; Wang et al. 2029u; Wang et al. 2029v; Wang et al. 2029w; Wang et al. 2029x; Wang et al. 2029y; Wang et al. 2029z; Wang et al. 2030a; Wang et al. 2030b; Wang et al. 2030c; Wang et al. 2030d; Wang et al. 2030e; Wang et al. 2030f; Wang et al. 2030g; Wang et al. 2030h; Wang et al. 2030i; Wang et al. 2030j; Wang et al. 2030k; Wang et al. 2030l; Wang et al. 2030m; Wang et al. 2030n; Wang et al. 2030o; Wang et al. 2030p; Wang et al. 2030q; Wang et al. 2030r; Wang et al. 2030s; Wang et al. 2030t; Wang et al. 2030u; Wang et al. 2030v; Wang et al. 2030w; Wang et al. 2030x; Wang et al. 2030y; Wang et al. 2030z; Wang et al. 2031a; Wang et al. 2031b; Wang et al. 2031c; Wang et al. 2031d; Wang et al. 2031e; Wang et al. 2031f; Wang et al. 2031g; Wang et al. 2031h; Wang et al. 2031i; Wang et al. 2031j; Wang et al. 2031k; Wang et al. 2031l; Wang et al. 2031m; Wang et

¹<https://github.com/ericniebler/NeuralCodeSum>

¹SEMERU / Computer Science Department, William and Mary, USA

[illegible][illegible][illegible]

BACKGROUND

VulRepair: A T5-Based Automated Software Vulnerability Repair

Michael Fu
ch.edu

Chakkrit Tantithamthavorn
chakrit@monash.edu
Monash University
Australia

Trung Le
trunglm@monash.edu
Monash University
Australia

Dinh Phung
dinh.phung@monash.edu
Monash University

Monash Software

Michael Fu
yeh.fu@monash.edu
Monash University
Australia

Chakkrit Tantithamthavorn*
chakkrit@monash.edu
Monash University
Australia

Trung Le
trunglm@monash.edu
Monash University
Australia

Dinh Phung
dinh.phung@monash.edu
Monash University
Australia

ing Conference and (ESEC/FSE '22), November 14–16, 2022, New York, NY, USA, 13 pages. <https://doi.org/10.1145/3541235>

ABSTRACT

ABSTRACT

As software vulnerabilities grow in volume and complexity, researchers proposed various Artificial Intelligence (AI)-based approaches to help under-resourced security analysts to find, detect, and localize vulnerabilities. However, security analysts still have to spend a huge amount of effort to manually fix or repair such Automated Vulnerability Repair, but it is still far from perfect due to various limitations. In this paper, we propose **VuREPAIR**, a TS-based automated software vulnerability repair to address various technical limitations of prior work. Through an extensive experiment with over 8,482 vulnerability fixes from real-world software projects, we find that our **VuREPAIR** achieves a better Prediction of 94%, which is 13%-21% more accurate than competitive baseline approaches. These results lead us to conclude that our **VuREPAIR** is considerably more accurate than two baseline approaches, highlighting the substantial advancement of NMT-based Automated Vulnerability Repairs. Our additional investigation also shows that our **VuREPAIR** can accurately repair as many as 745 out of 1,706 real-world, well-known vulnerabilities (e.g., Use After Free, Improper Input Validation, OS Command Injection), demonstrating the practicality and significance of our **VuREPAIR** for generating vulnerability repairs, helping under-resourced security analysts on fixing vulnerabilities.

INTRODUCTION

Software vulnerabilities found in under-... may r... to att... extor... the in... the

SECURITY

ing - Security and privacy:

CCS CONCEPTS

Software and... Nguyen, and Dinh...

KEYWORDS

KEYWORDS
Software Vulnerability Repair
Reference Format: [Tantitham](#)

Software Vulnerability Reference Format:

ACM Reference Format:
Michael Fu, Chakkrit Tantithamthornchai, and Phung. 2022. VulRepair: A T5-Based Automatic Repair Tool for C/C++ Code. In *Proceedings of the ACM Conference on Computer Security*. ACM, New York, NY, 12 pages.
https://doi.org/10.1145/3512345.3512346

The corresponding author

Permission to make digital or hard copies of this work for personal or classroom use is granted without fee provided that copies are not made for profit or commercial advantage and that copies bear this notice and the first page on the first page. Copyrights for components of this work owned by others than Springer Nature must be honored. Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. This article is published with open access at <https://doi.org/10.1007/s11227-022-04511-1>.
© The Author(s) 2022. Singapore, Singapore
14-18, 2022, Singapore, Singapore
Machinery

1 INTRODUCTION

1 INTRODUCTION

Software vulnerabilities are security flaws, glitches, or weaknesses found in software systems that could be exploited by attackers to undertake malicious activities [6]. In particular, criminal groups may make use of unrescued security vulnerabilities in software to attack and damage a system to steal confidential information (NVD) or extort assets, resulting in severe economic damage [17]. According to the statistics of the National Vulnerability Database (NVD), the number of vulnerabilities discovered per year in 2021 [1], increased five times from 4k/year in 2011 to 20k+/year in 2021 [1].

Recently, researchers proposed various Artificial Intelligence (AI)-based approaches to help under-resourced security analysts better understand the characteristics of vulnerabilities (e.g., vulnerability analysis) [2, 3, 5, 7, 44, 57, 58, 62], find vulnerabilities (e.g., vulnerability prediction) [12, 20, 21, 25, 26, 34, 37, 43, 49], and assess the impact of vulnerabilities (e.g., vulnerability prediction or better analysis) [2, 3, 5, 7, 44, 57, 58, 62].

For example, the AI predicts if a given function in [37, 43, 65],

ability analysis approaches are proposed to predict the level of given function (e.g., function [12], syntactic [13], semantic [20, 43], syntactic [20, 21, 28, 36, 38–40, 48, 59]) using various types of information (e.g., graph properties [12, 21, 55, 65], and mutual information [21, 28, 36, 38–40, 48, 59]) and mutual information only help security analysts to find, detect, and localize the location of vulnerabilities. However, security analysts still have to spend a huge amount of effort to manually fix or repair vulnerabilities [4, 11, 35].

Recently, Chen et al. [13] proposed VRepair [13], an automated vulnerability repair approach that leverages a Transformer-based Neural Machine Translation (NMT) in the area of automated program repairs (e.g., see Seacrest [14], an RNN-based NMT model). However, VRepair is still inaccurate due to the following three limitations:

- VRepair is trained on a small bug fix corpus of 100 functions, which may generate suboptimal level tokenization (OOV).

Limitation Q: VRepair is trained on 23,607 C/C++ functions, which may generate various bugs. VRepair is trained on various (e.g., Sequence) features. VRepair is not accurate on all bugs.

Limitation R: VRepair leverages the word-level tokenization and the copy mechanism to handle Out-Of-Vocabulary (OOV), limiting its ability to generate new tokens that never appear in a vulnerable function but are newly introduced in the vulnerability repair.

Identifier-aware Unified Pre-trained Encoder-Decoder Model for Code Understanding and Generation

¹ Salesforce Research Asia
² Nanyang Technological University, Singapore
y, weishi.wang, sjoty, shoi}@salesforce.com

which can be transferred to benefit downstream tasks, especially those with little notation. Inspired by their success, recent attempts to adapt these pre-trained models for programming language (PL) tasks have achieved promising results on code

However, despite their success, models rely on either an encoder-only (Svyatkovskiy et al., 2020) or a decoder-only model (Beltagy et al., 2020), which is suboptimal for reading and understanding tasks, respectively. CodeBERT (Feng et al., 2020) is a pre-trained model that

additional decoder when a
marization task, where the
fit from the pre-training.
methods simply employ the
training techniques on so
as a sequence of tokens
nores the rich structural
-mitted to fully compre

In this work, we propose an encoder-decoder model that takes type information in the T5 architecture and employs denoising self-supervised pre-training and supervised training for understanding and generation. In addition, we assign developer-assigned

ing programs, distinctive identifiers are available, so that they serve rich code identifiers in rationality. To we propose trains the identifiers.

Further

-
-
-

An Empirical Study on the Usage of Transformer Models for Code Completion

Abstract—Code completion aims at speeding up code writing by suggesting the next line of code in this field focused on improving the productivity of the programmer. In this paper, we present a new code completion system, called *Code Completion*, which is able to suggest the next line of code in a program. The system is based on a machine learning approach, and it is able to learn from a large number of programs. The system is able to suggest the next line of code in a program, and it is able to learn from a large number of programs. The system is able to suggest the next line of code in a program, and it is able to learn from a large number of programs.

[illegible]

PRODUCTION

Completion is considered as one of the "killer" features of modern Integrated Development Environments [63, 64, 72]. It can provide developers with information about the next code to be developed, and

completions considering the context surrounding the code [63, 72], the history of code changes [72], and/or code patterns mined from software [63, 64, 72].

[illegible][illegible]

Studying the Usage of Text-To-Text Transfer to Support Code-Related Tasks

Department, William and Mary, USA

For example, in the work by Watson *et al.* [116] the input string representing a test method without an assert state and the output is an appropriate assert statement for given test. In the approach by Haque *et al.* [117], the input is composed of strings representing a subtest to document while the output is a natural language summary document the subtest.

Recent years have seen the rise of transfer learning in the field of natural language processing. The basic idea is to pre-train a model on a large and generic dataset by using supervised task, e.g., making tokens in strings and using the model to guess the masked tokens. Then, the trained model is fine-tuned on smaller and specialized datasets, each one at supporting a specific task. In this context, Raffel *et al.* proposed the T5 (Text-To-Text Transfer Transformer) model, pre-trained on a large natural language corpus and fine-tuned to achieve state-of-the-art performance on many tasks characterized by text-to-text transformation.

The goal of this work is to empirically investigate the potential of a TS model when pre-trained and fine-tuned to support many of the previously listed code-related tasks, characterized by text-to-text transformations. We start by pre-training a TS model using a large dataset consisting of 499,618 English sentences and 1,569,889 source code comments (i.e., methods). Then, we fine-tune the model using datasets from previous work with the goal of supporting code-related tasks:

- **Automatic bug-fixing.** We use the dataset by Tufano et al. [10], composed of instances in which the “input” string is transformed by a human fixer method and the “target”

is the fixed version of the same method.

Injection of code snippets. This dataset is also by [Tobin et al. \[15\]](#) and features instances in which the input strings are reversed as compared to automatic bug-fixing: the input is a fixed method, while the output is its fixed version). The model must learn here to inject bugs (snippets) in code instead of fixing bugs.

Generation of assert statements in test methods. We used dataset by [Watson et al. \[16\]](#), composed of instances in which the input string is a representation of a test method with an assert statement and a focal method it tests (i.e., the

production method tested), while the output string encodes appropriate assert statement for the input test method.

--	--

021

2021

Learning to Recommend Method Names with Glo

Fang Liu
Key Lab of High Confidence Software
Technology, MoE (Peking University)
Beijing, China
liufang816@pku.edu.cn

Ge Li^{*}
Key Lab of High Confidence Software
Technology, MoE (Peking University)
Beijing, China
lige@pku.edu.cn

Yiyang Hao
Silicon Heart Tech Co., Ltd
Beijing, China
haoyiyang@nthink.com

Shuai Lu
Key Lab of High Confidence Software
Technology, MoE (Peking University)
Beijing, China
lushuai96@pku.edu.cn

Key Lab of H
Technolo

Key Lab
Techno

ABSTRACT

In programming, the names for the program entities, especially for the methods, are the intuitive characteristic for understanding the functionality of the code. To ensure the readability and maintainability of the programs, method names should be consistent with other names used in related contexts in their codebase. In recent years, many automated approaches are proposed to suggest consistent names for methods, among which neural machine translation (NMT) based models are widely used and have achieved state-of-the-art results. However, these NMT-based models mainly focus on extracting the code-specific features from the method body or the surrounding methods, the project-specific features are ignored. We conduct a statistical analysis to explore the relationship between the method name and its contexts. Based on the statistical results, the project-specific context, which considers the local context, the project-specific context, and the documentation of the method simultaneously. Experimental results on Java methods show that our model can outperform the state-of-the-art results by a large margin on method name suggestion, demonstrating the effectiveness of our proposed model.

CCS CONCEPTS

• Software and its engineering: • Computing methodologies
• Artificial intelligence.

KEYWORDS

method name recommendation, global context, deep learning

^{*}Corresponding authors.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, to republish, to post on servers or to redistribute to lists, requires prior specific permission and/or fee. Request permissions from permissions@acm.org.

JCSSE '22, May 21–29, 2022, Pittsburgh, PA, USA
© 2022 Association for Computing Machinery.
ACM 9781450392211-1/22/05...\$15.00
JCSSE '22, May 21–29, 2022, Pittsburgh, PA, USA
© 2022 Association for Computing Machinery.
ACM 9781450392211-1/22/05...\$15.00

InferFix: End-to-End Program Repair with LLMs over Retrieval-Augmented Prompts

Matthew Jin
Microsoft
Redmond, WA, USA

Xin Shi
Microsoft
Redmond, WA, USA

Syed Shahriar
UCLA
Los Angeles, CA, USA

Shuai Lu
Microsoft Research
Beijing, China

Alexey Svyatkovskiy
Microsoft
Redmond, WA, USA

Michele T
Micro
Redmond

Neel S
M
Redm

ABSTRACT

Software development life cycle is profoundly influenced by bugs; their introduction, identification, and eventual resolution account for a significant portion of software development cost. This has motivated software engineering researchers and practitioners to propose different approaches for automating the identification and repair of software defects.

Large language models have been adapted to the program repair task through few-shot demonstration learning and instruction prompting, treating this as an infilling task. However, these models have only focused on learning general bug-fixing patterns for uncategorized bugs mined from public repositories. In this paper, we propose InferFix: a transformer-based program repair framework paired with a state-of-the-art static analyzer to fix critical security and performance bugs. InferFix combines a Retriever—transformer encoder model pretrained via contrastive learning objective, which aims at searching for semantically equivalent bugs and corresponding fixes; and a Generator—a large language model (12 billion parameter Code Cushman model) finetuned on supervised bug-fix data with prompts augmented via adding bug type annotations and semantically similar fixes retrieved from an external non-parametric memory.

To train and evaluate our approach, we curated InferFixBugs, a novel, metadata-rich dataset of bugs extracted by executing the Infer static analyzer on the change histories of thousands of Java and C# repositories. Our evaluation demonstrates that InferFix outperforms strong LLM baselines, with a top-1 accuracy of 65.6% for generating fixes in C# and 76.8% in Java. We discuss the deployment of InferFix alongside Infer at Microsoft which offers an end-to-end solution for detection, classification, and localization of bugs, as well as fixing and validation of candidate patches, integrated in the continuous integration pipeline to automate the software development workflow.

KEYWORDS

Program repair, static analyses, prompt augmentation, finetuning

ACM Reference Format:

Matthew Jin, Syed Shahriar, Michele T, Neel S, Alexey Svyatkovskiy, Shuai Lu, Xin Shi, Fang Liu, Ge Li, Zhiyi Fu, Shuangyue Burgh, PA, USA, ACM, New York, NY, 2022. 13 Mar 2023

1 INTRODUCTION

The software development process is a critical component of the software development process. In well-tested projects, most code files are paired with at least one corresponding test file that implements unit and/or integration test functions that evaluate the functionality of the code. However, writing high quality tests can be time-consuming [1, 2] and is often either partially or entirely neglected. This has led to extensive work in automated test generation, including both black-box [3, 4, 5, 6] and white-box [7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25, 26, 27, 28, 29, 30, 31, 32, 33, 34, 35, 36, 37, 38, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49, 50, 51, 52, 53, 54, 55, 56, 57, 58, 59, 60, 61, 62, 63, 64, 65, 66, 67, 68, 69, 70, 71, 72, 73, 74, 75, 76, 77, 78, 79, 80, 81, 82, 83, 84, 85, 86, 87, 88, 89, 90, 91, 92, 93, 94, 95, 96, 97, 98, 99, 100] approaches.

Classical test generation tools like EvoSuite [4] directly optimize to generate high-coverage tests. However, the generated tests are often hard to read and may be unrealistic or even verify generated test correctness [5].

^{*}Equal contribution

Software testing is a critical component of the software development process. In well-tested projects, most code files are paired with at least one corresponding test file that implements unit and/or integration test functions that evaluate the functionality of the code. However, writing high quality tests can be time-consuming [1, 2] and is often either partially or entirely neglected. This has led to extensive work in automated test generation, including both black-box [3, 4, 5, 6] and white-box [7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25, 26, 27, 28, 29, 30, 31, 32, 33, 34, 35, 36, 37, 38, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49, 50, 51, 52, 53, 54, 55, 56, 57, 58, 59, 60, 61, 62, 63, 64, 65, 66, 67, 68, 69, 70, 71, 72, 73, 74, 75, 76, 77, 78, 79, 80, 81, 82, 83, 84, 85, 86, 87, 88, 89, 90, 91, 92, 93, 94, 95, 96, 97, 98, 99, 100] approaches.

Training Language Models for Code And Tests

Nikhitha Rao^{*}, Kush Jain^{*}, Uri Alon, Claire Le Goues, Vincent J. Hellendoorn, Nikitha Rao, kdjain, urialon, legoues, vhellendoorn}@cmu.edu

Abstract—Testing is an integral but often neglected part of the software development process. Classical test generation tools such as EvoSuite generate behavioral test suites by optimizing for coverage, but tend to produce tests that are less well-suited for the highly similar to that written by humans, as current models are trained to generate code that can generate code that is natural language processing, and thus fail to consider the under-test context when producing a test file. In this work, we propose the Aligned Code And Tests Language Model (CAT-LM), a novel pretraining signal that explicitly considers the code context is available to the model when generating test code. We analyze its usefulness for realistic applications, showing that it efficiently produces tests that achieve coverage similar to ones written by developers while resembling their writing style. By utilizing the code context, CAT-LM generates more valid tests than even much larger language models trained with more data (CodeGen 16B and StarCoder) and substantially outperforms a recent test-specific model (TeCo) at test completion. Overall, our work highlights the importance of incorporating software-specific insights when training language models for code and test generation. **Index Terms**—test generation, test completion, large language models, code-test alignment

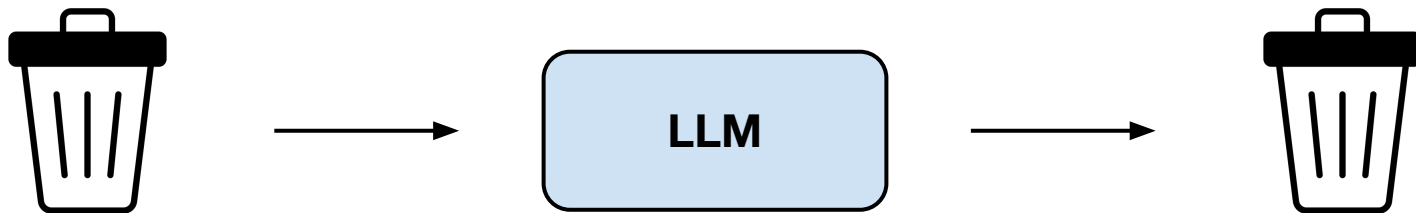
1. INTRODUCTION

Software testing is a critical component of the software development process. In well-tested projects, most code files are paired with at least one corresponding test file that implements unit and/or integration test functions that evaluate the functionality of the code. However, writing high quality tests can be time-consuming [1, 2] and is often either partially or entirely neglected. This has led to extensive work in automated test generation, including both black-box [3, 4, 5, 6] and white-box [7, 8, 9, 10, 11, 12, 13, 14, 15, 16, 17, 18, 19, 20, 21, 22, 23, 24, 25, 26, 27, 28, 29, 30, 31, 32, 33, 34, 35, 36, 37, 38, 39, 40, 41, 42, 43, 44, 45, 46, 47, 48, 49, 50, 51, 52, 53, 54, 55, 56, 57, 58, 59, 60, 61, 62, 63, 64, 65, 66, 67, 68, 69, 70, 71, 72, 73, 74, 75, 76, 77, 78, 79, 80, 81, 82, 83, 84, 85, 86, 87, 88, 89, 90, 91, 92, 93, 94, 95, 96, 97, 98, 99, 100] approaches.

Classical test generation tools like EvoSuite [4] directly optimize to generate high-coverage tests. However, the generated tests are often hard to read and may be unrealistic or even verify generated test correctness [5].

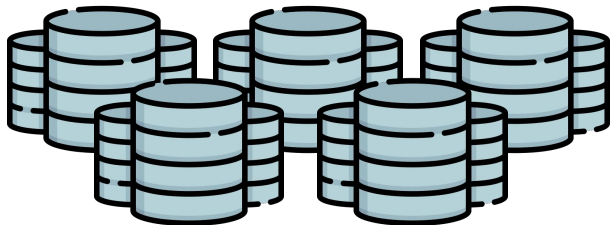
MOTIVATION

Large Language Models (LLMs) are data-driven techniques, so “**Garbage in, garbage out**” (GIGO) remains actual.



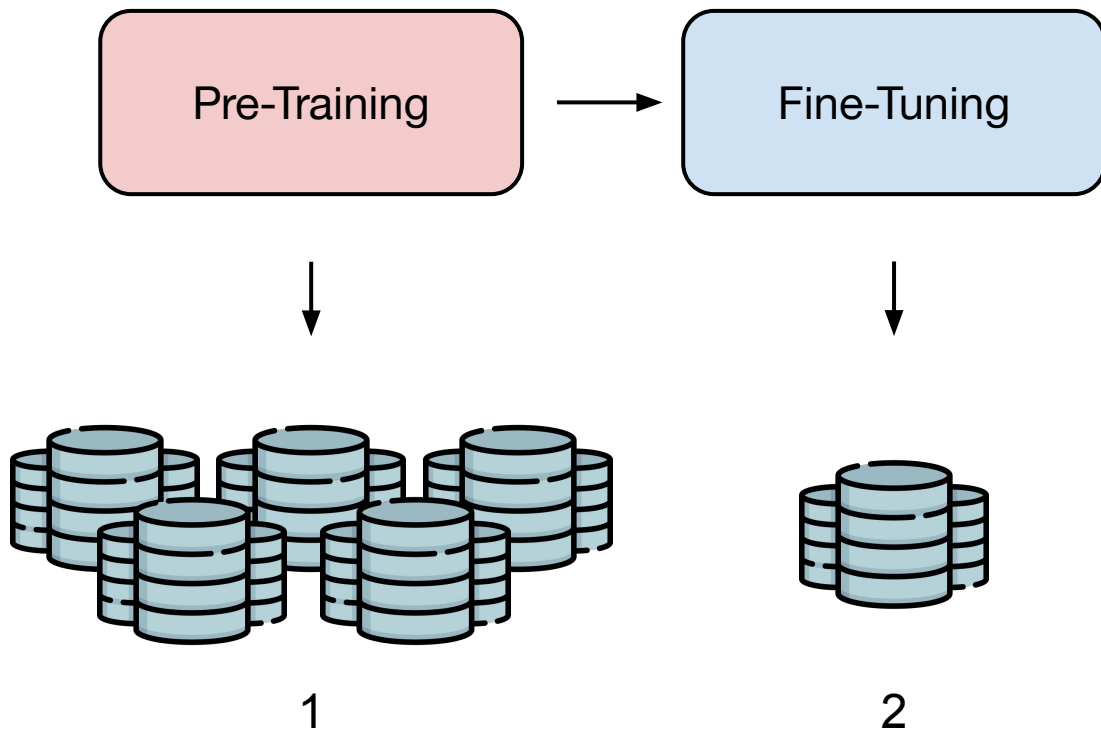
DATA IMPORTANCE

Pre-Training

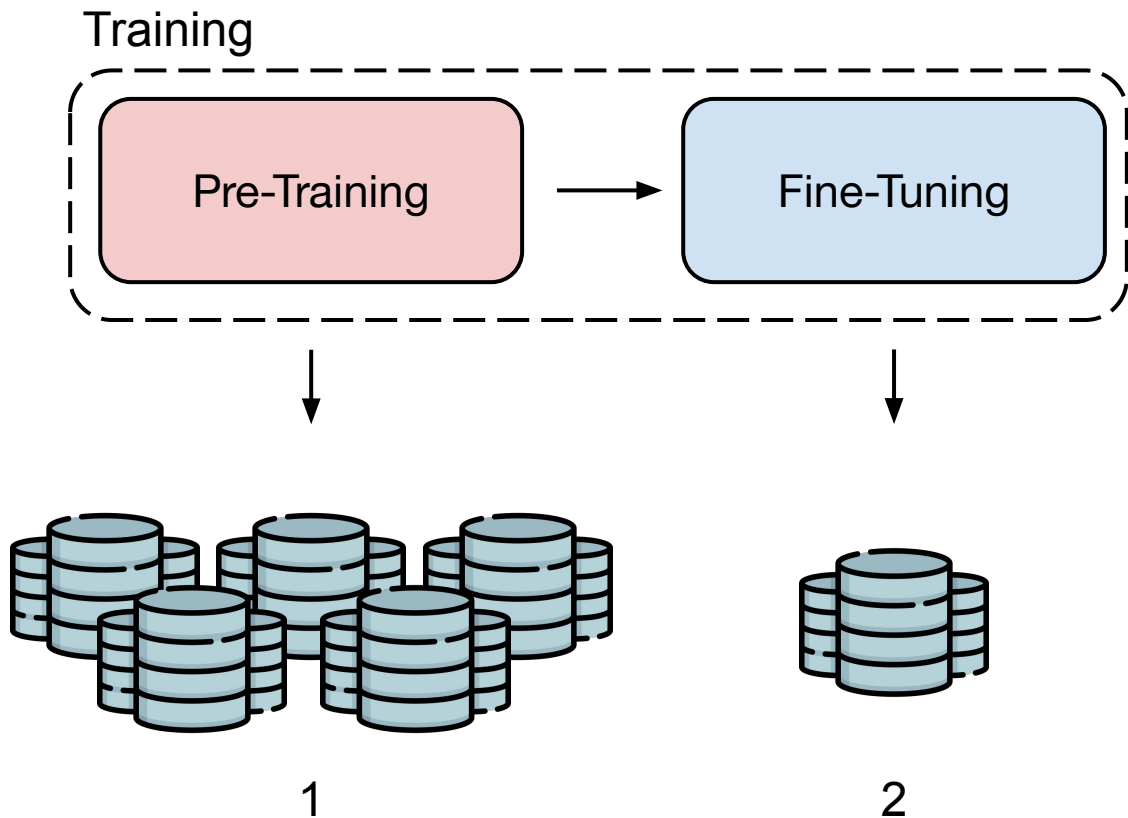


1

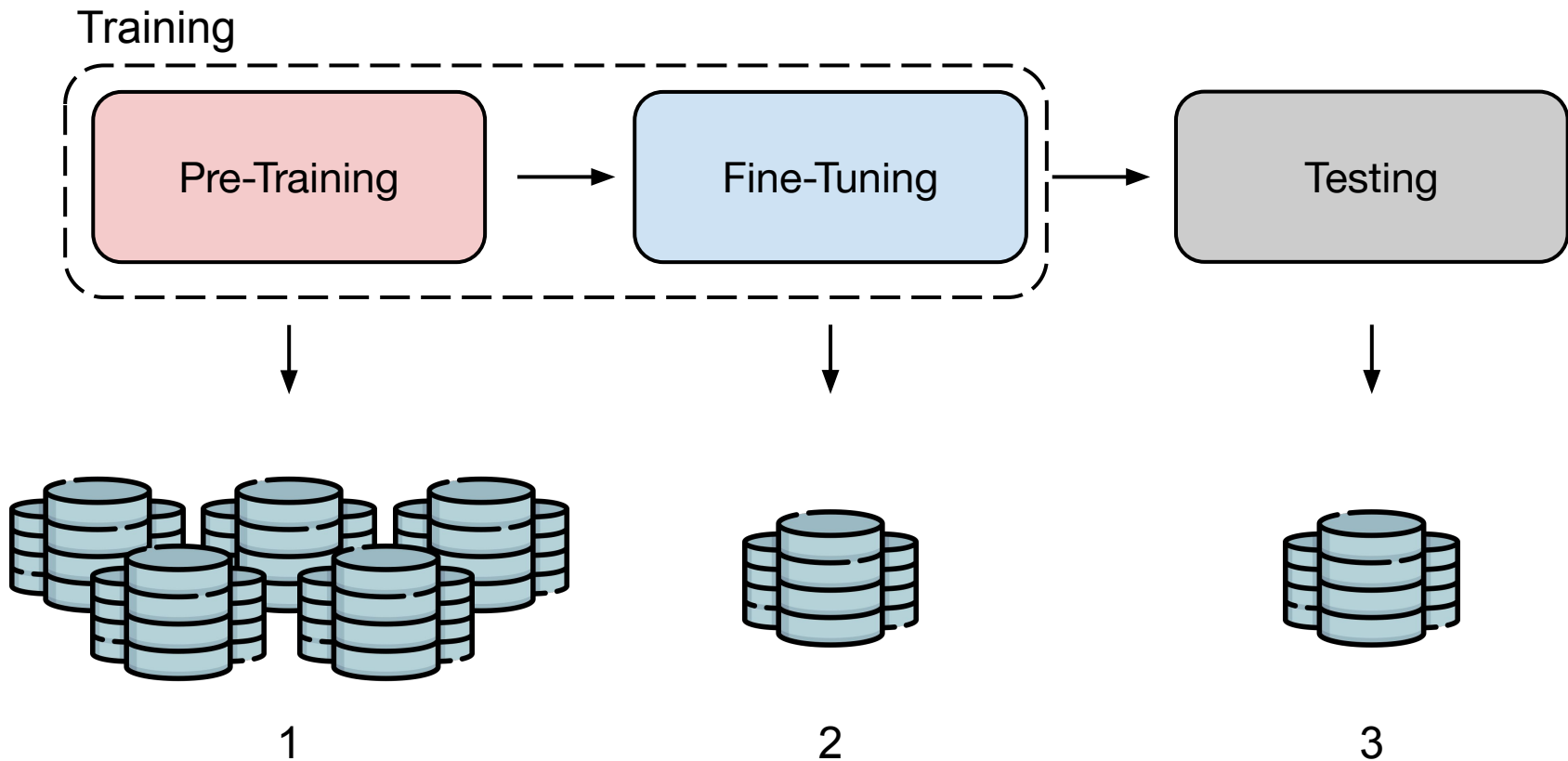
DATA IMPORTANCE



DATA IMPORTANCE



DATA IMPORTANCE



DATA IMPORTANCE



**LOW-QUALITY
INSTANCES**

**HIGH-QUALITY
INSTANCES**

DATA IMPORTANCE

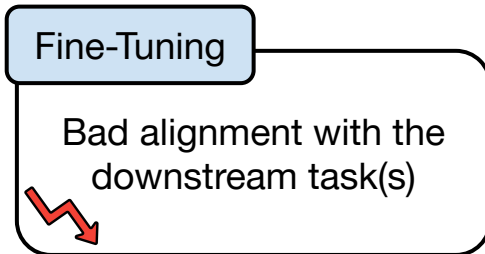
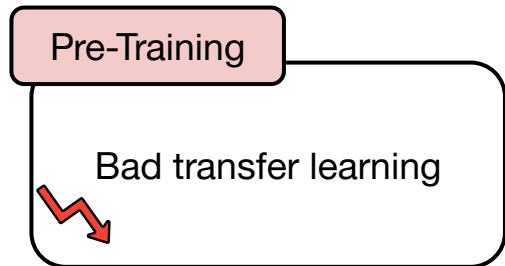


Pre-Training

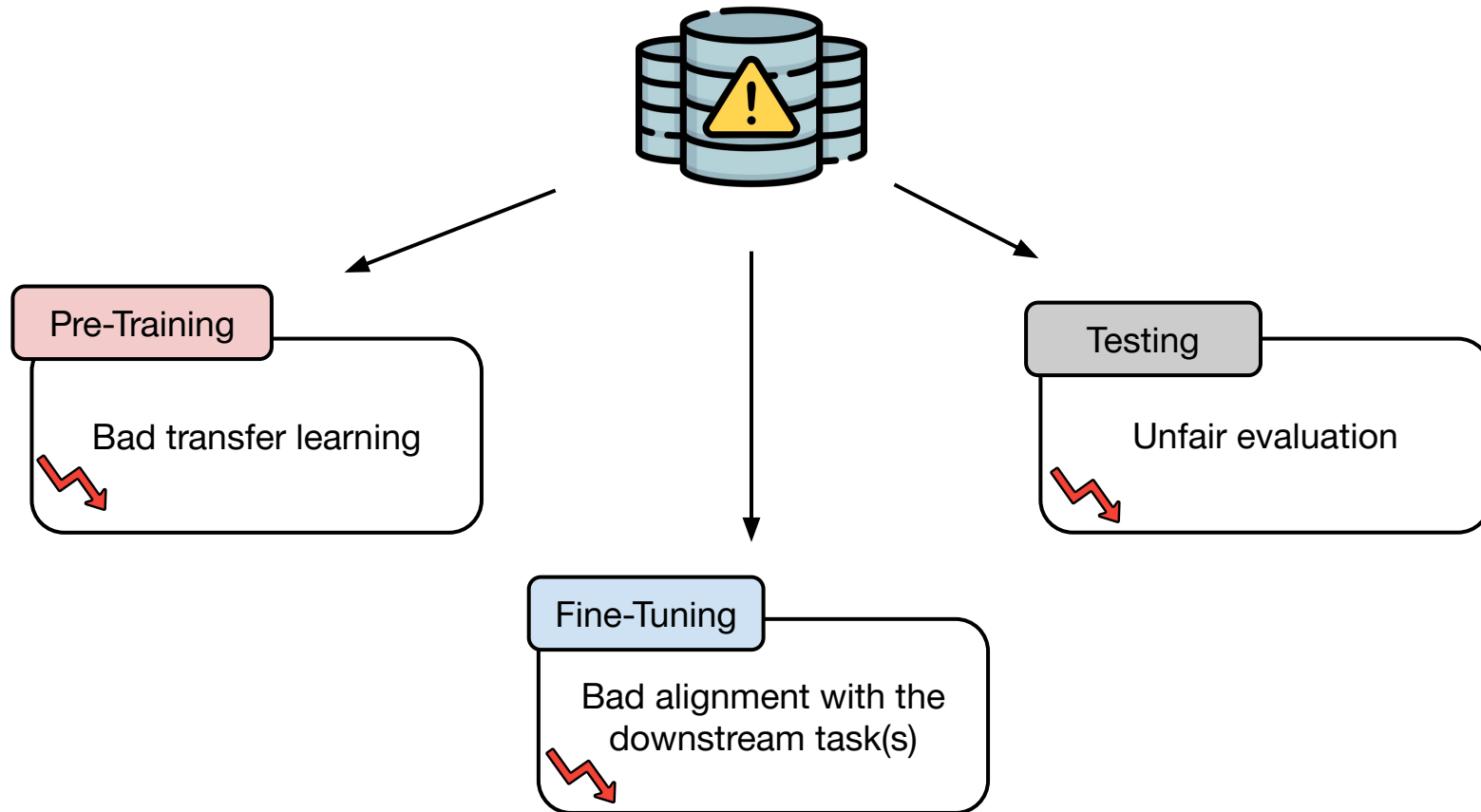
Bad transfer learning



DATA IMPORTANCE



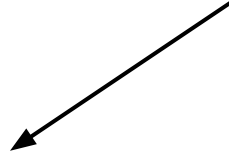
DATA IMPORTANCE



DATA IMPORTANCE

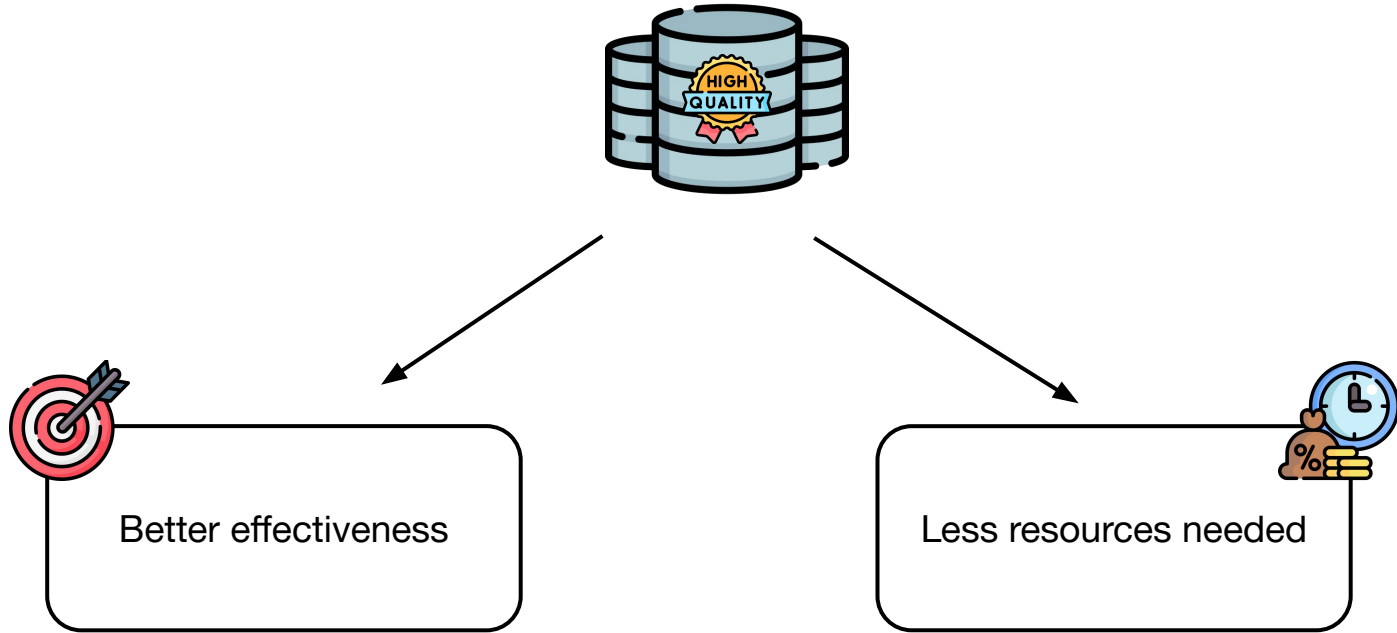


DATA IMPORTANCE



Better effectiveness

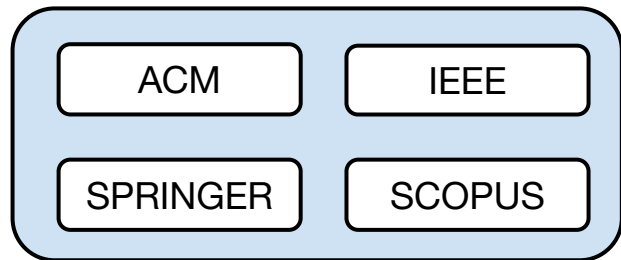
DATA IMPORTANCE



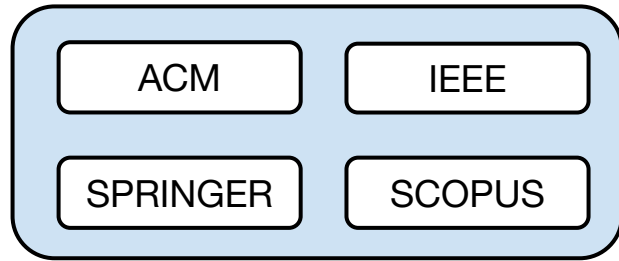
What are the **known data smells** for code-related tasks and the **strategies** adopted to remove them?



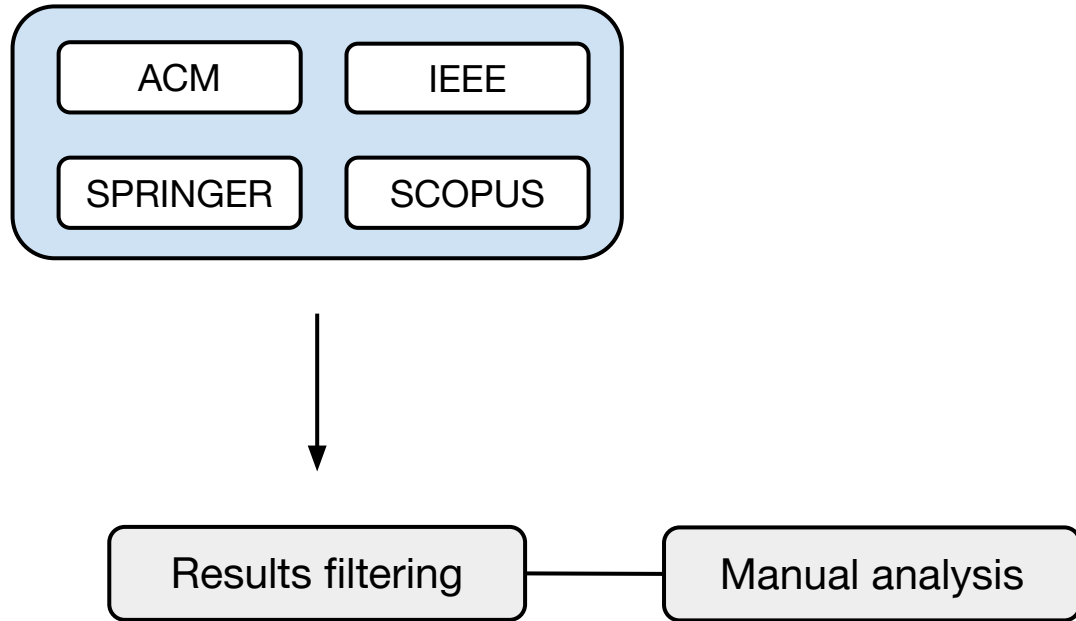
DESIGN



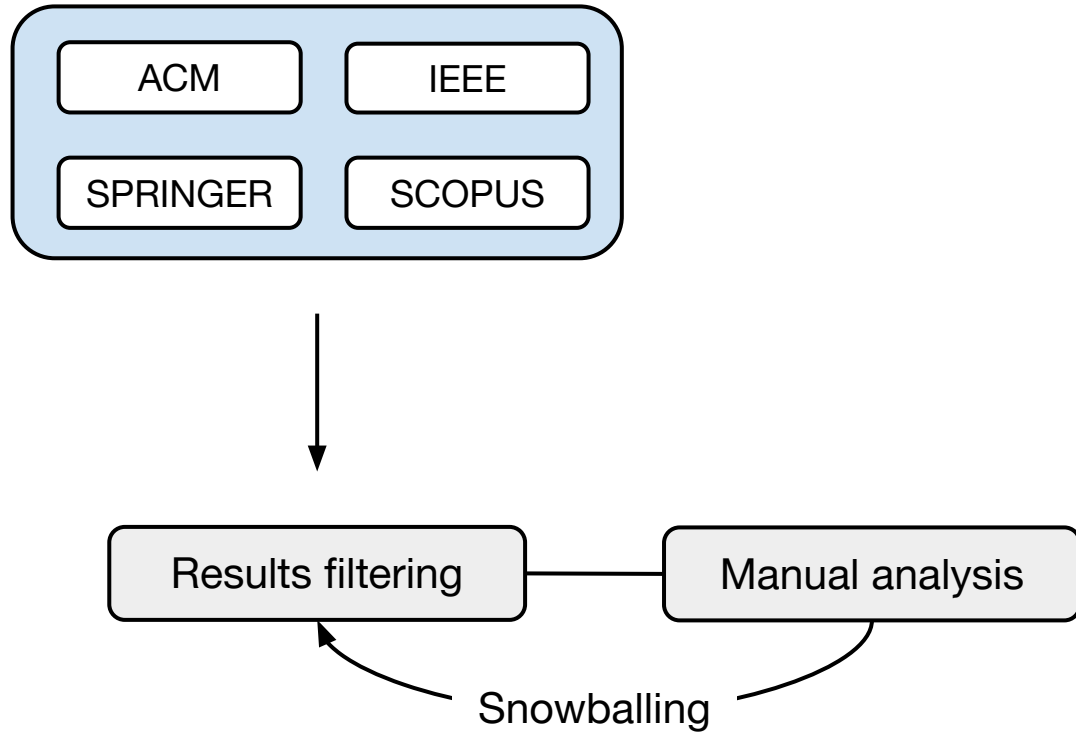
DESIGN

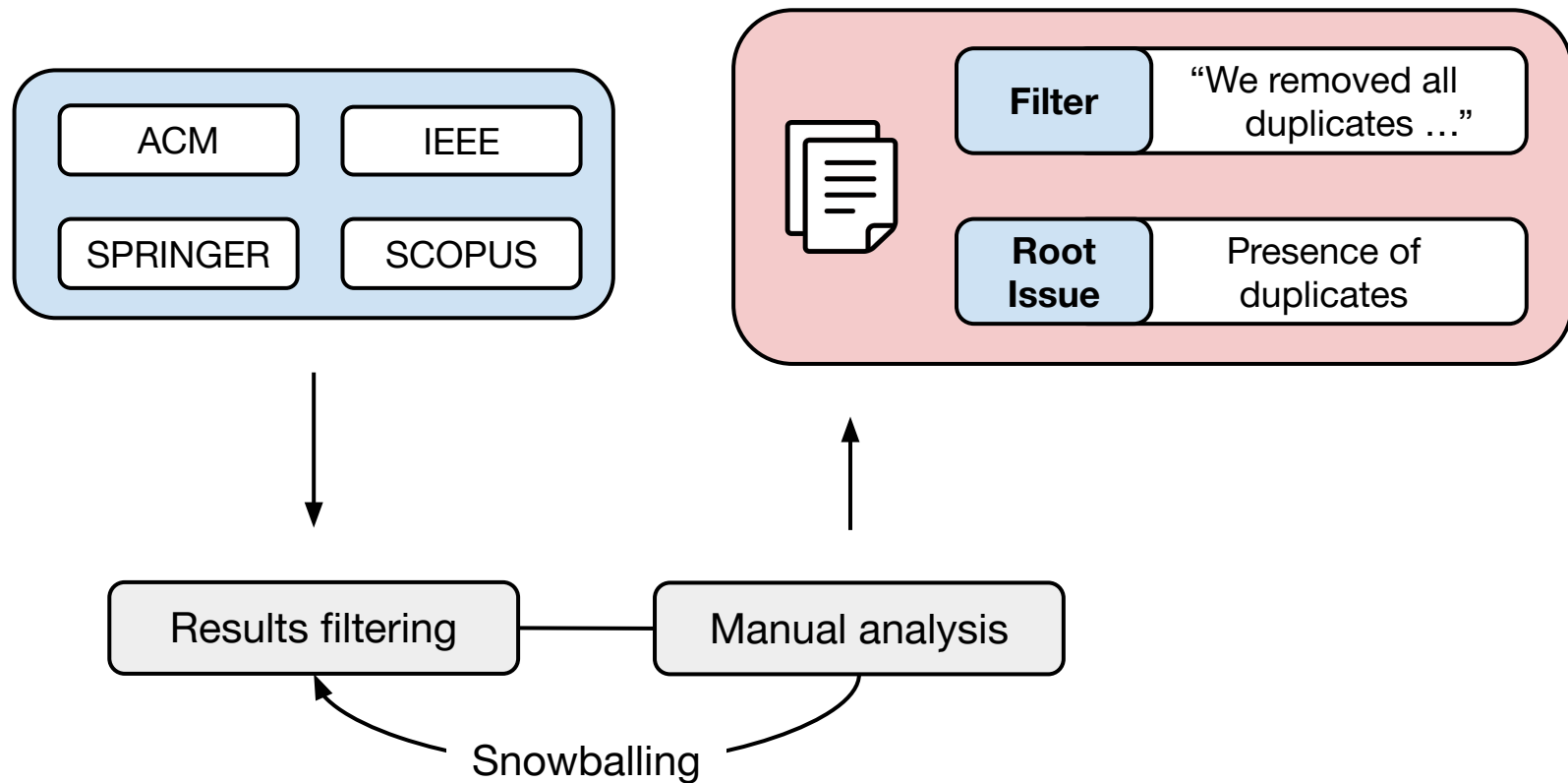


DESIGN

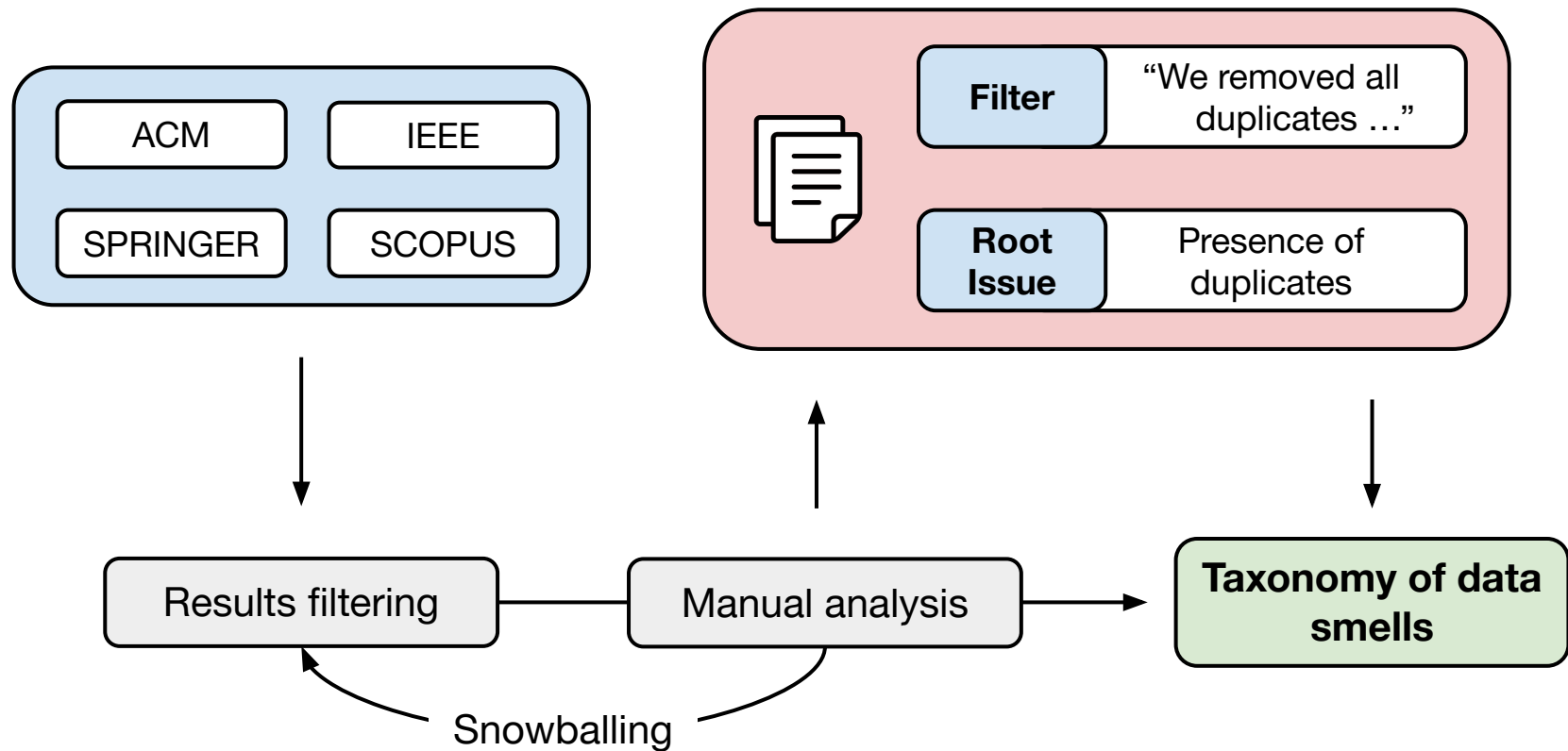


DESIGN

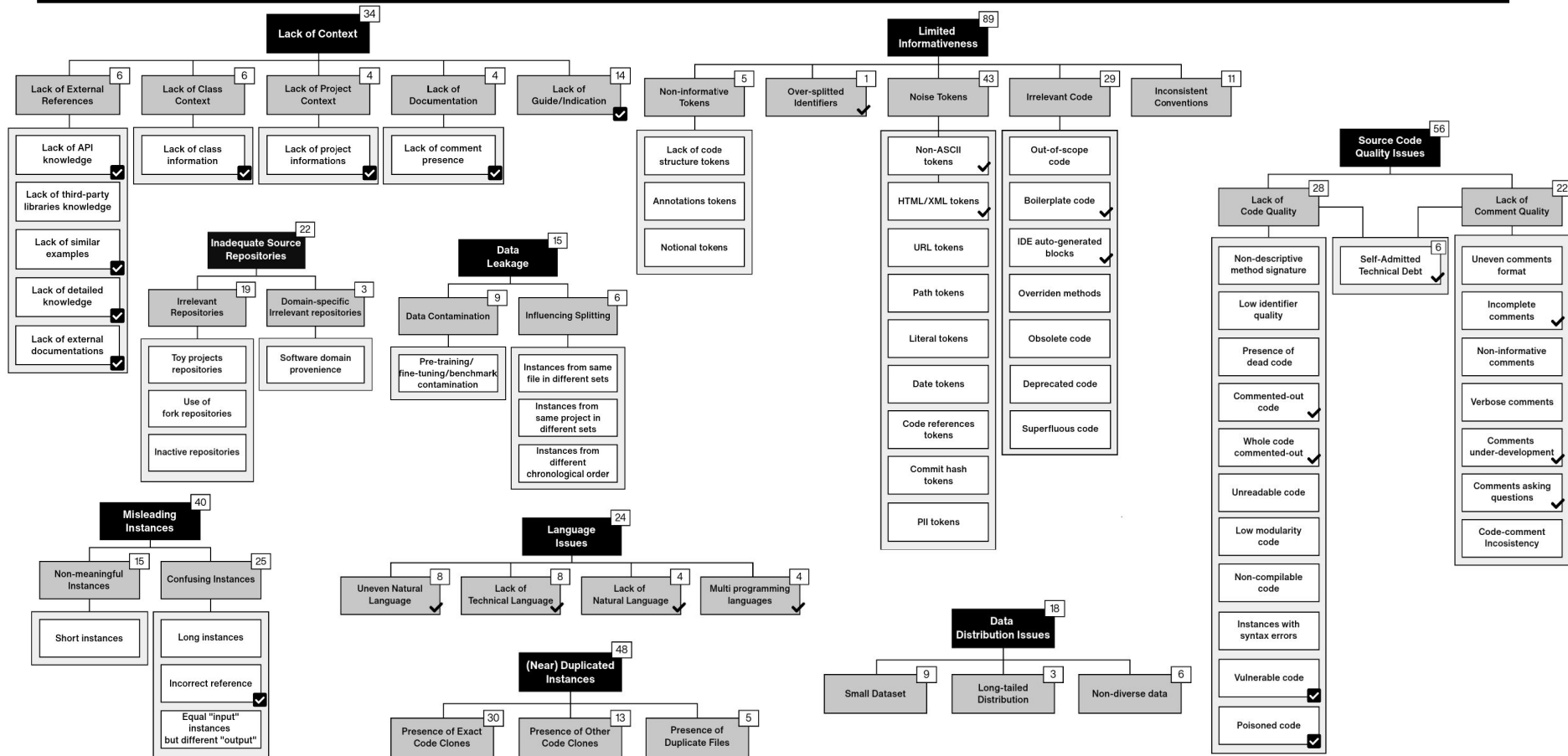




DESIGN



RQ1



RQ1

Limited Informativeness

```
issue_title = "Fix NPE in UserService 🐛"

issue_body = """
When calling getUserById() it throws NullPointerException ☹️.

Stacktrace:
java.lang.NullPointerException at com.app.service.UserService.java:42

Related commit: 9f3a1c7b2d4e8a91
Reported by: alice.dev@company.com

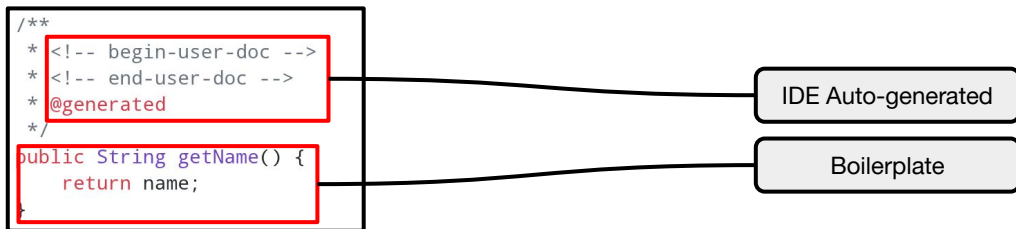
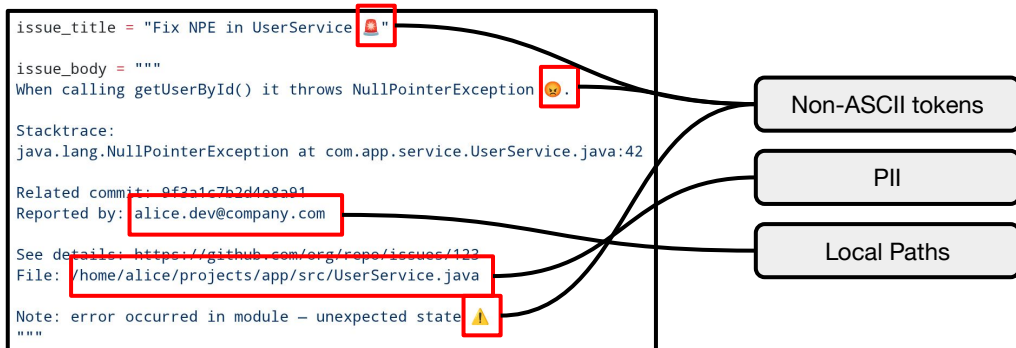
See details: https://github.com/org/repo/issues/123
File: /home/alice/projects/app/src/UserService.java

Note: error occurred in module – unexpected state ⚠️
"""
```

```
/**
 * <!-- begin-user-doc -->
 * <!-- end-user-doc -->
 * @generated
 */
public String getName() {
    return name;
}
```

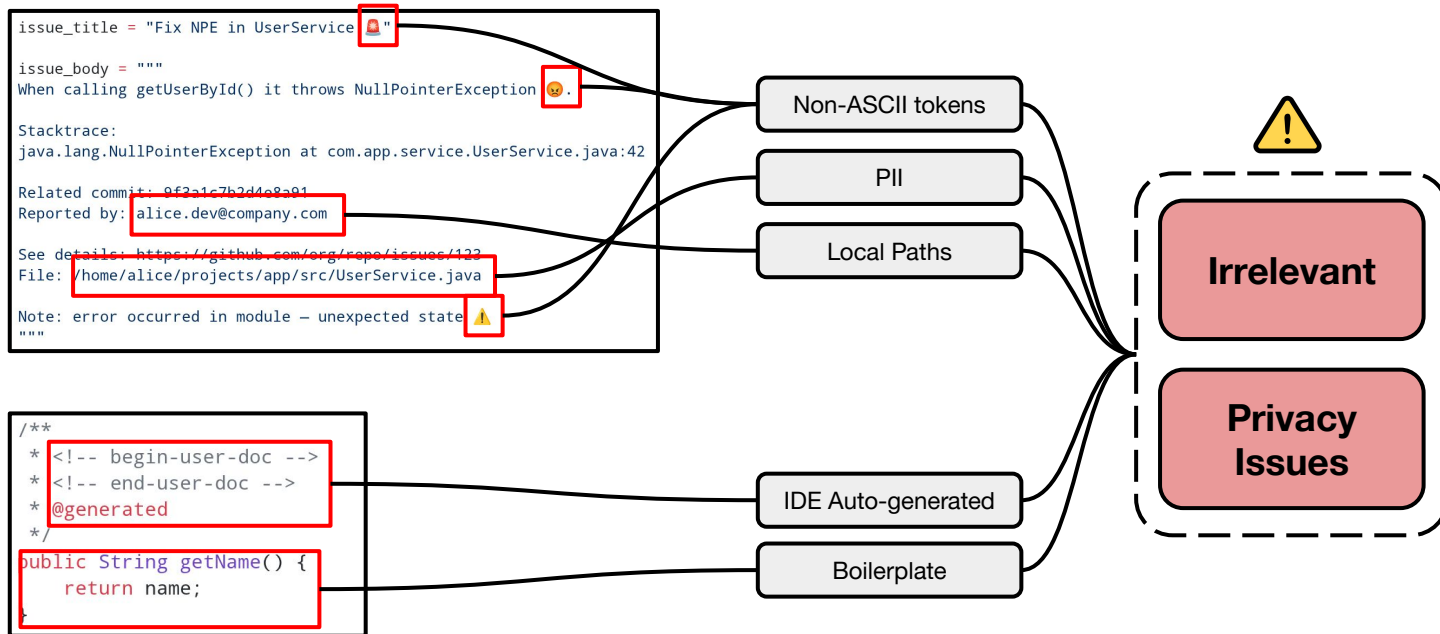
RQ1

Limited Informativeness



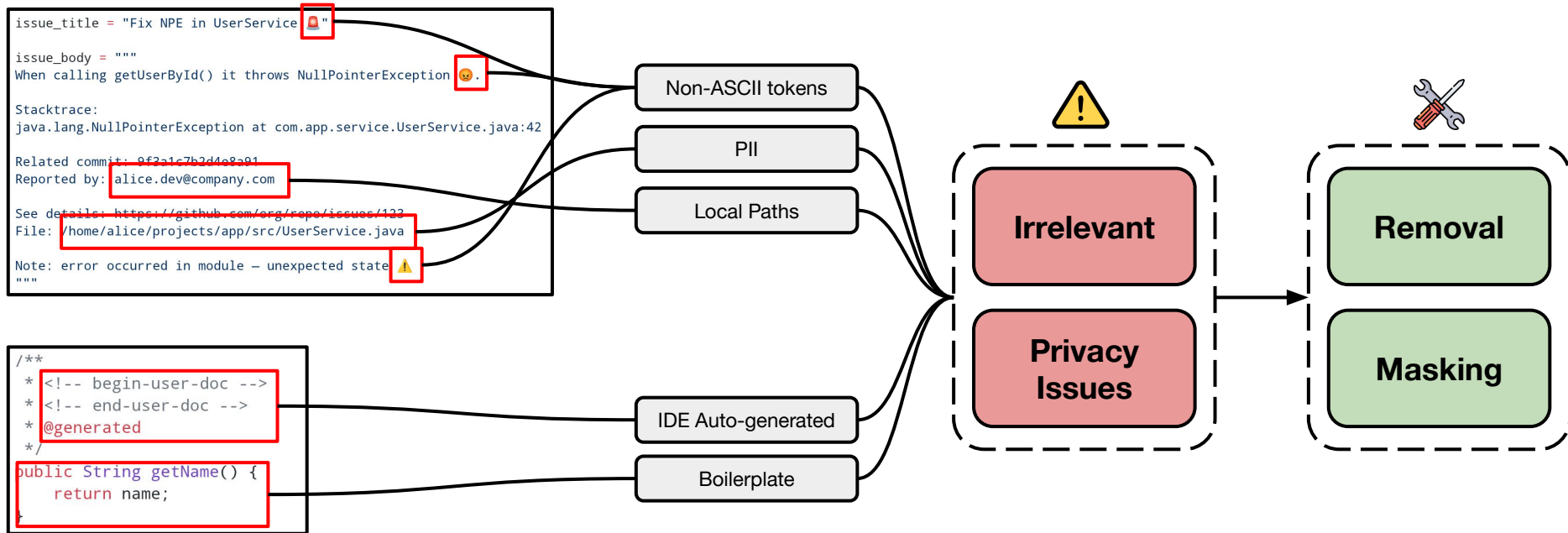
RQ1

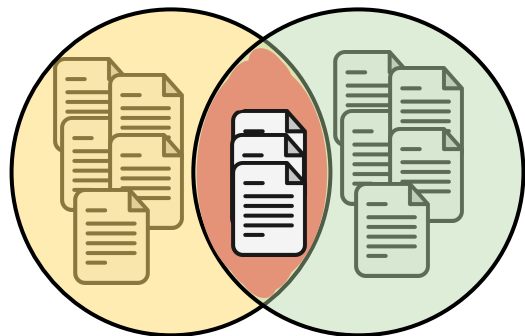
Limited Informativeness



RQ1

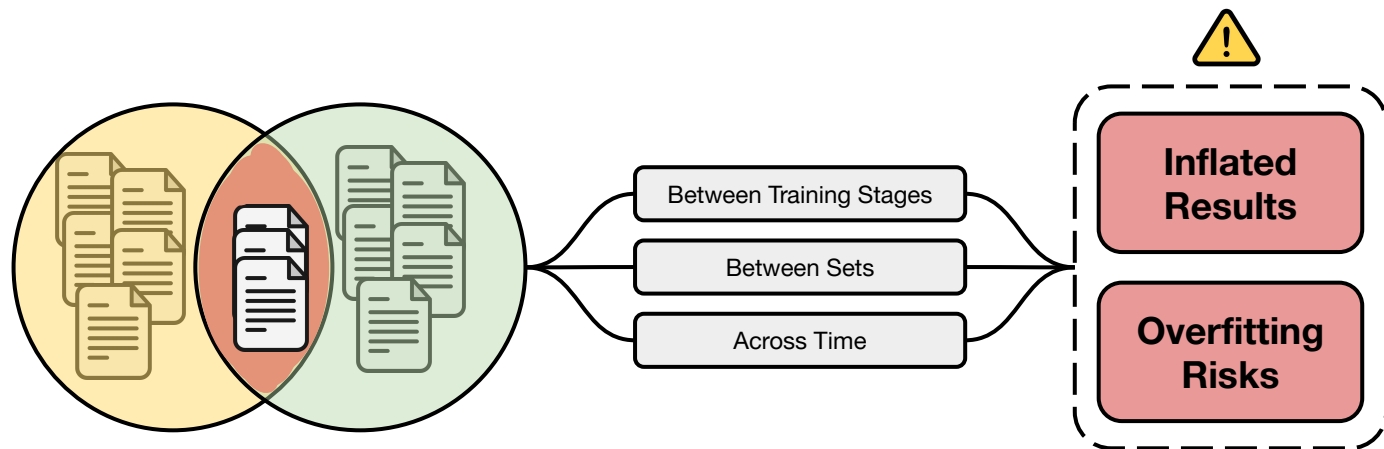
Limited Informativeness





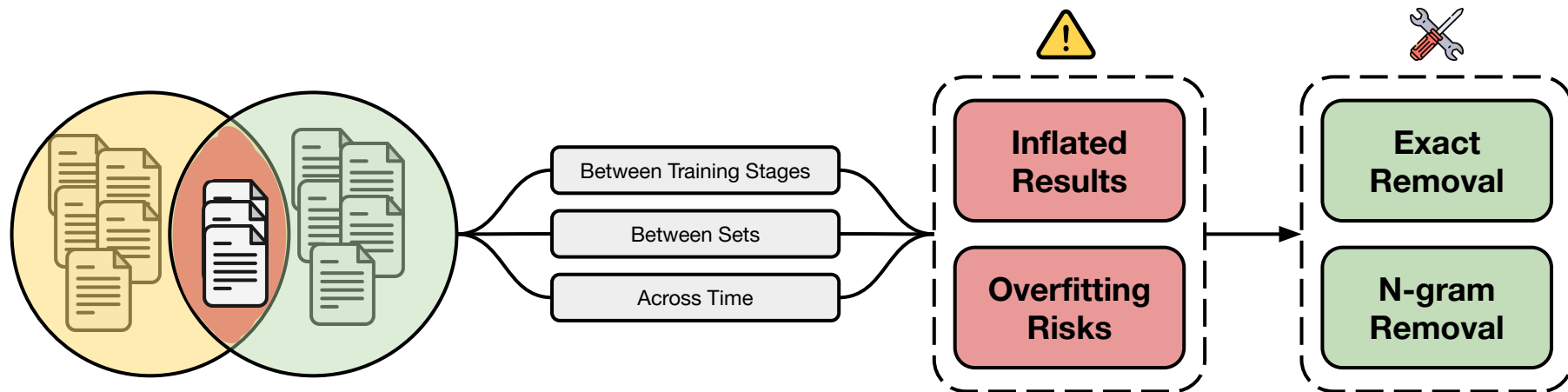
RQ1

Data Leakage



RQ1

Data Leakage



RQ1

Lack of Context

```
from typing import List
from pydantic import BaseModel
import spacy, networkx as nx, re, uuid

class Result(BaseModel):
    id: str
    categories: List[str]
    confidence: float

class ContextAnalyzer:
    def __init__(self):
        self.nlp = spacy.load("en_core_web_sm")
        self.graph = nx.DiGraph()

    def analyze(self, text: str) -> Result:
        doc = self.nlp(text)
        cats = self._detect(text, doc)

        r = Result(
            id=str(uuid.uuid4()),
            categories=cats,
            confidence=max(1 - 0.2 * len(cats), 0.1),
        )

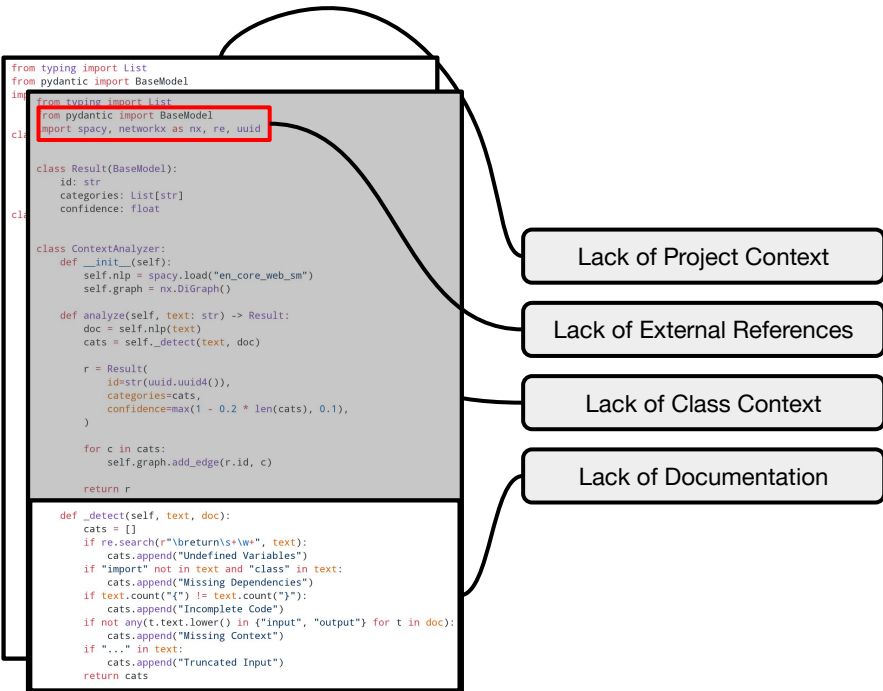
        for c in cats:
            self.graph.add_edge(r.id, c)

        return r

    def _detect(self, text, doc):
        cats = []
        if re.search(r"\breturn\b", text):
            cats.append("Undefined Variables")
        if "import" not in text and "class" in text:
            cats.append("Missing Dependencies")
        if text.count("{") != text.count("}"):
            cats.append("Incomplete Code")
        if not any(t.text.lower() in {"input", "output"} for t in doc):
            cats.append("Missing Context")
        if "..." in text:
            cats.append("Truncated Input")
        return cats
```

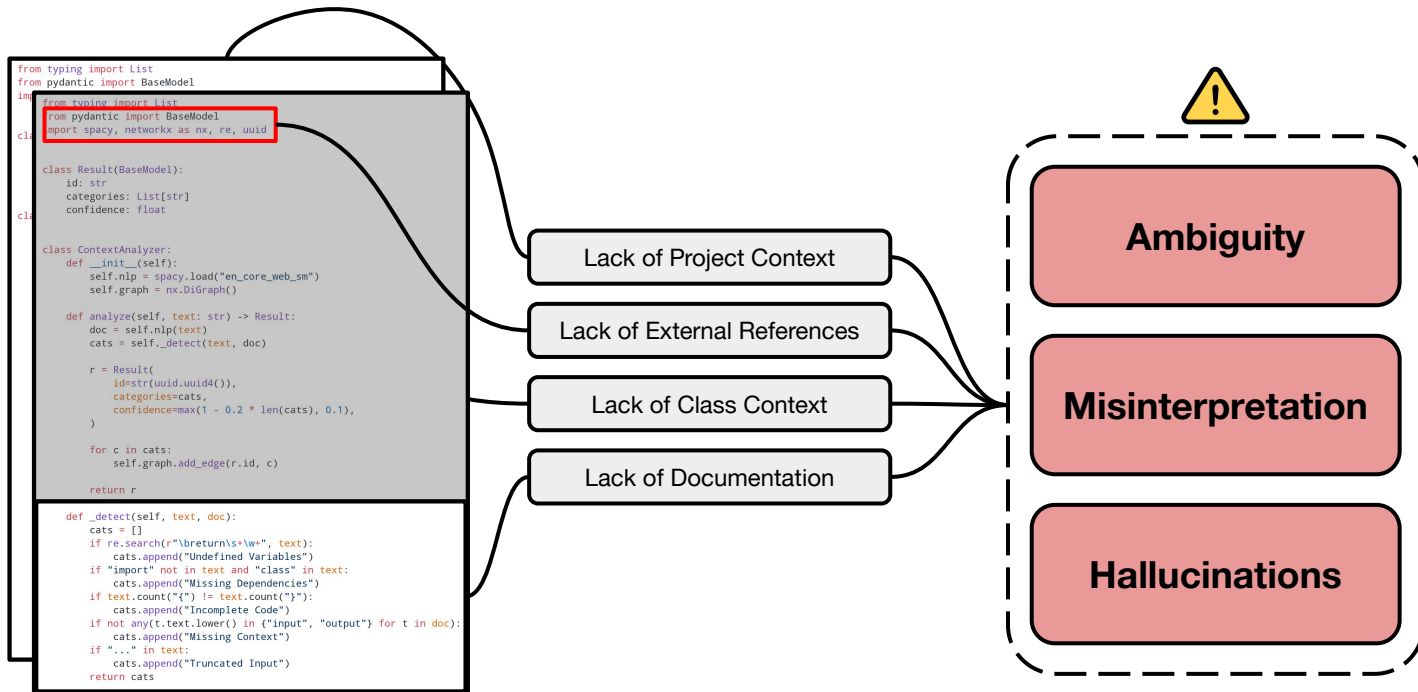
RQ1

Lack of Context



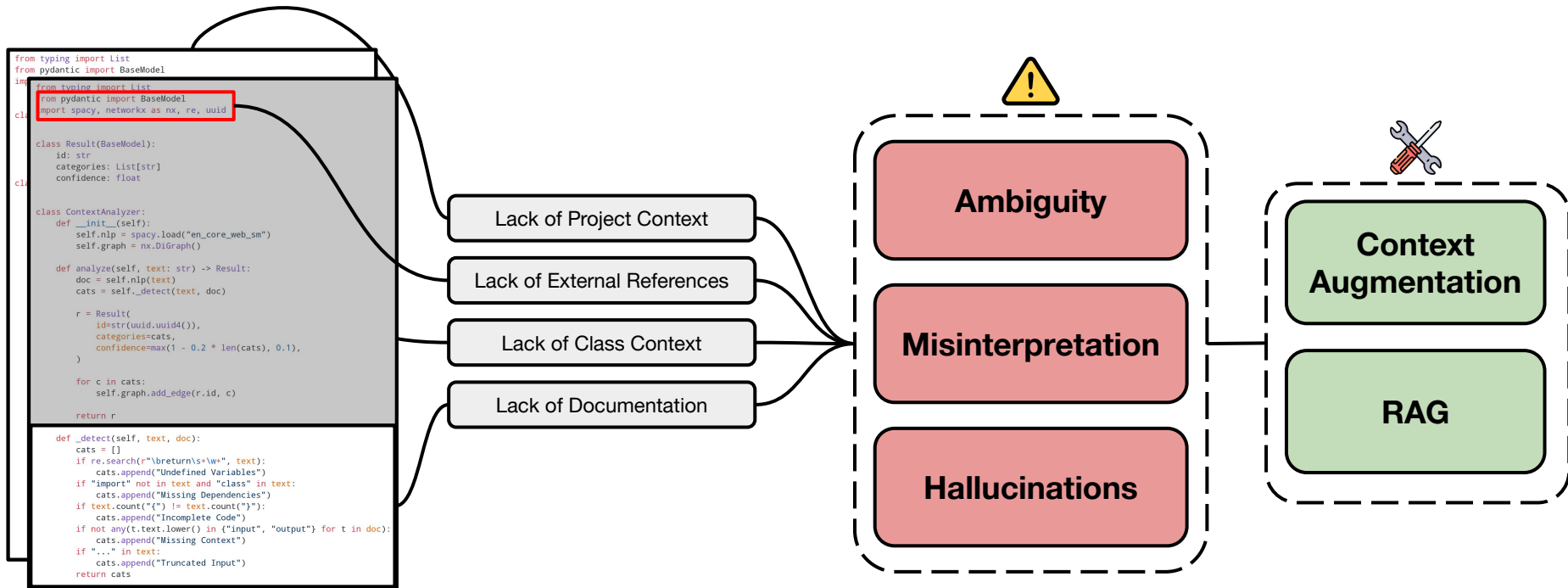
RQ1

Lack of Context



RQ1

Lack of Context



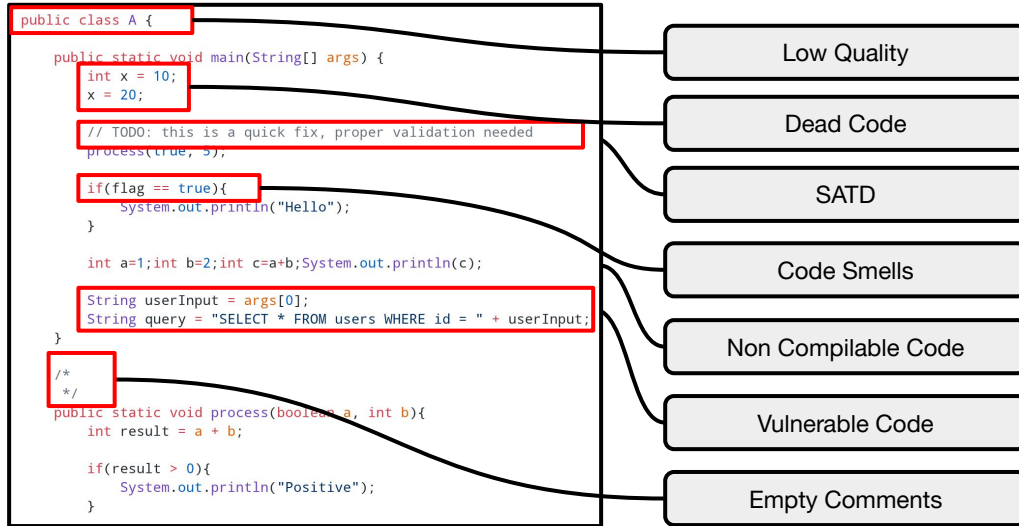
RQ1

Source Code Quality Issues

```
public class A {  
  
    public static void main(String[] args) {  
        int x = 10;  
        x = 20;  
  
        // TODO: this is a quick fix, proper validation needed  
        process(true, 5);  
  
        if(flag == true){  
            System.out.println("Hello");  
        }  
  
        int a=1;int b=2;int c=a+b;System.out.println(c);  
  
        String userInput = args[0];  
        String query = "SELECT * FROM users WHERE id = " + userInput;  
    }  
  
    /*  
    */  
    public static void process(boolean a, int b){  
        int result = a + b;  
  
        if(result > 0){  
            System.out.println("Positive");  
        }  
    }  
}
```

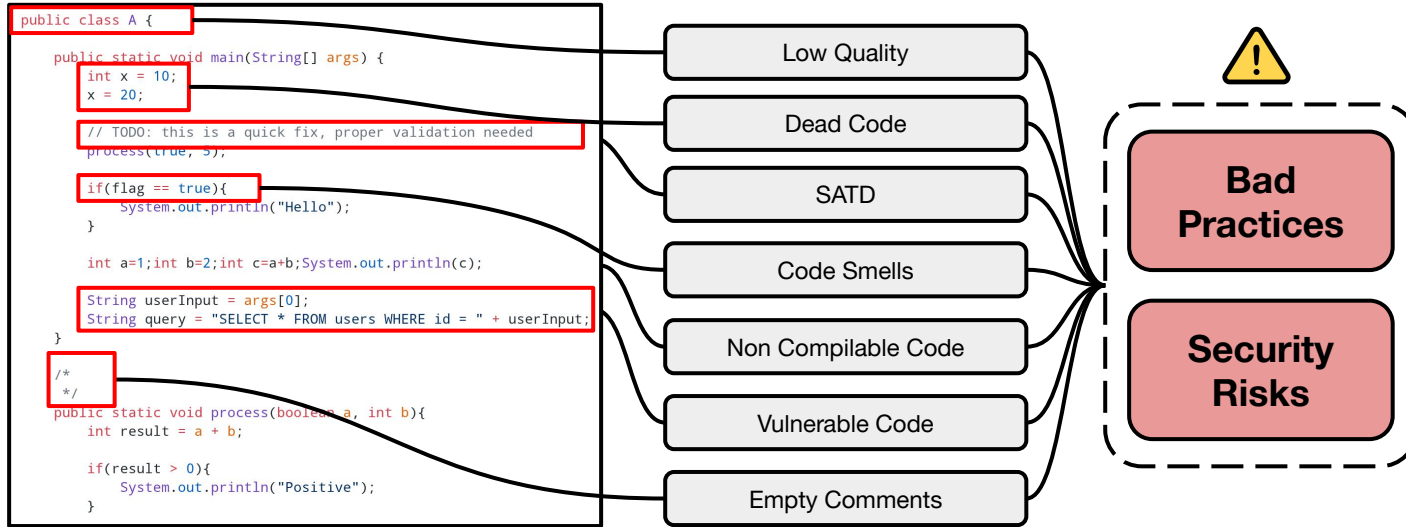
RQ1

Source Code Quality Issues



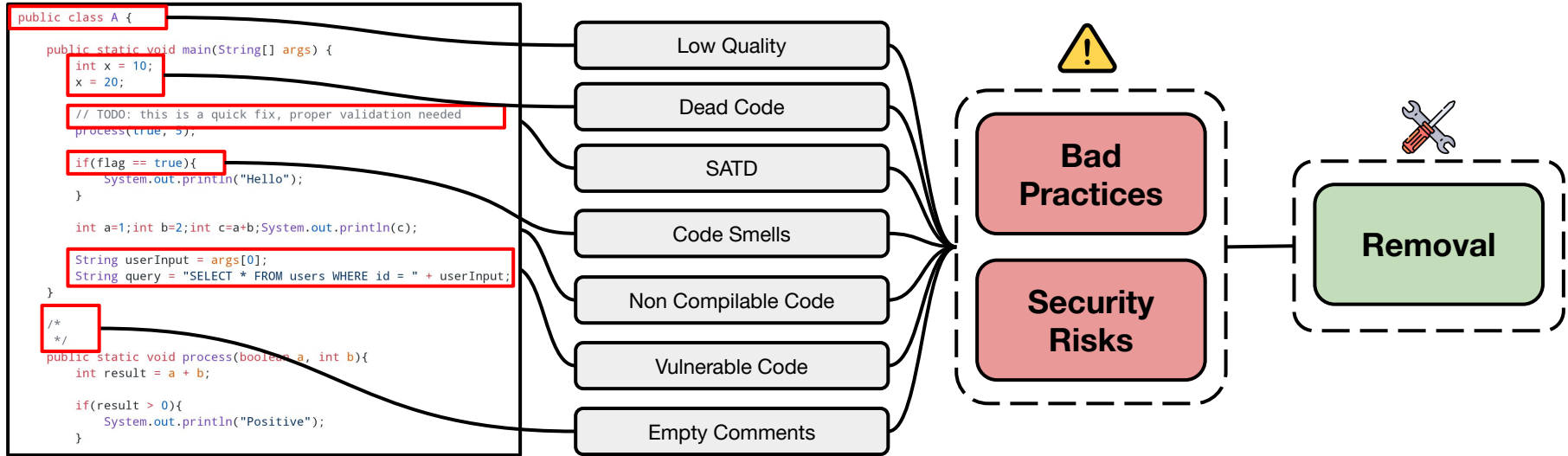
RQ1

Source Code Quality Issues



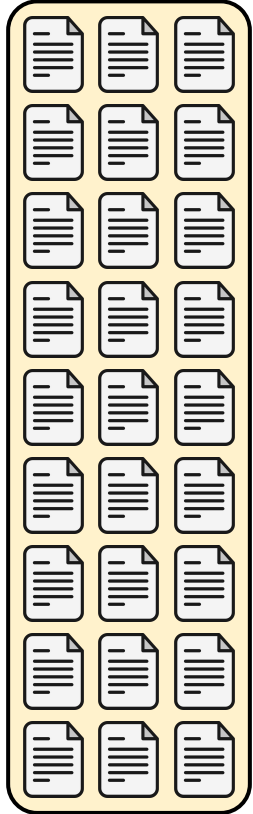
RQ1

Source Code Quality Issues



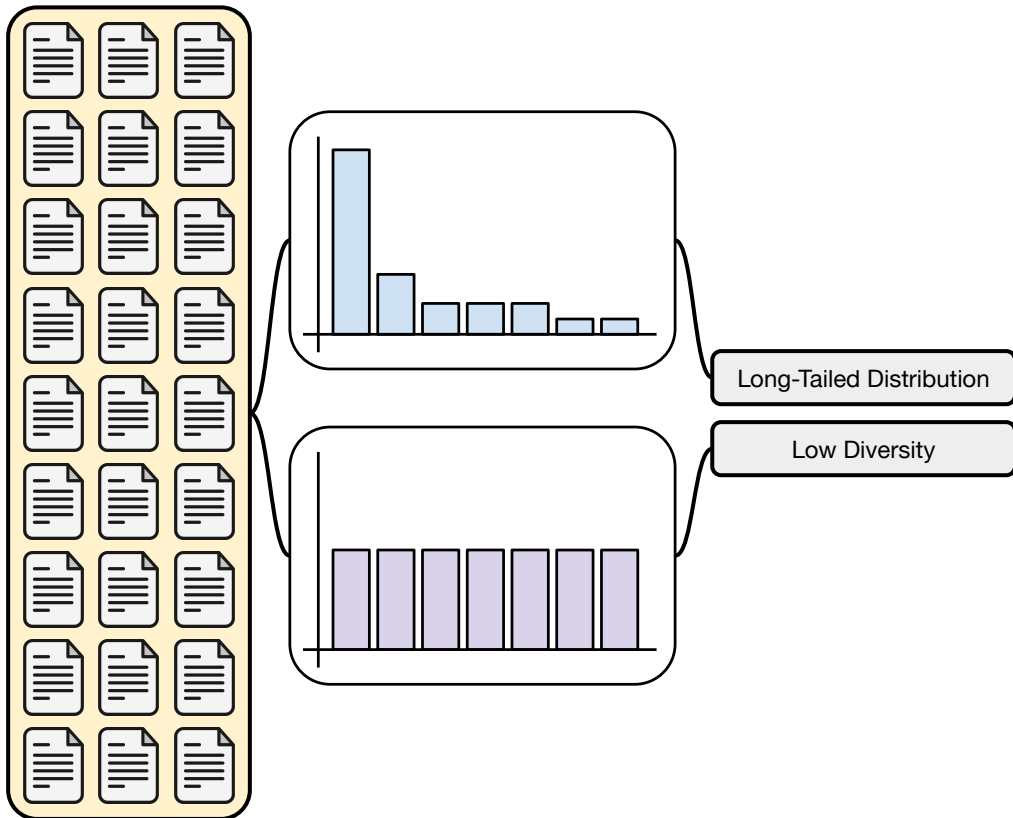
RQ1

Data Distribution Issues



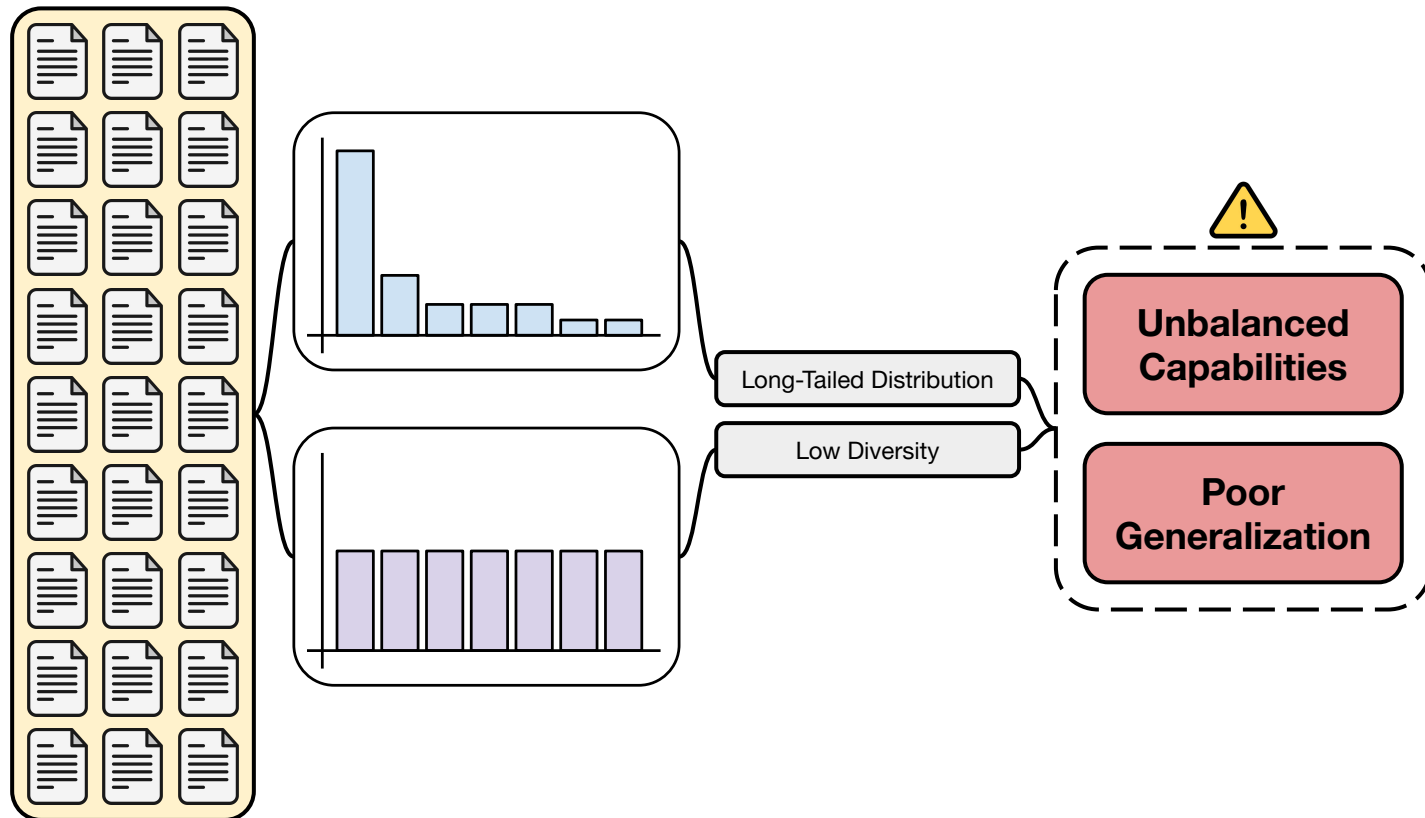
RQ1

Data Distribution Issues



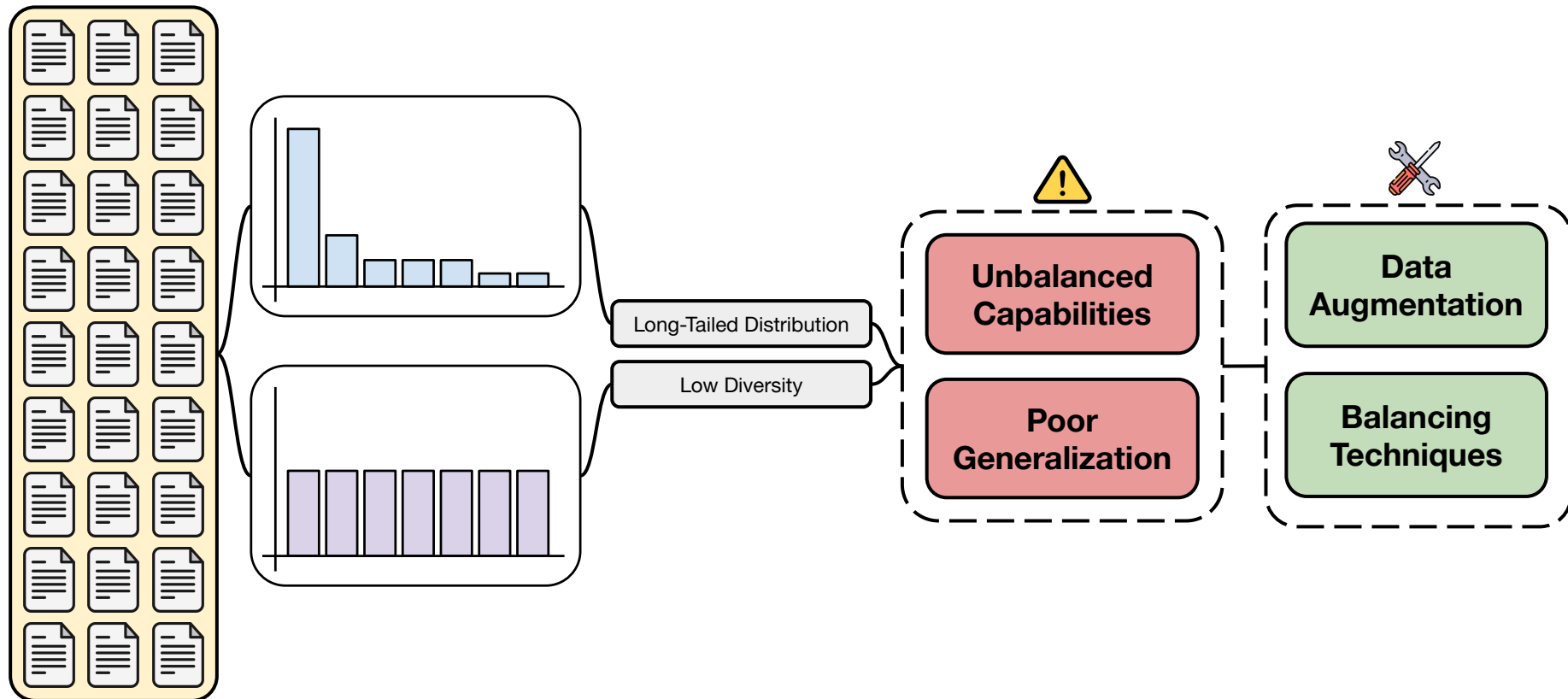
RQ1

Data Distribution Issues



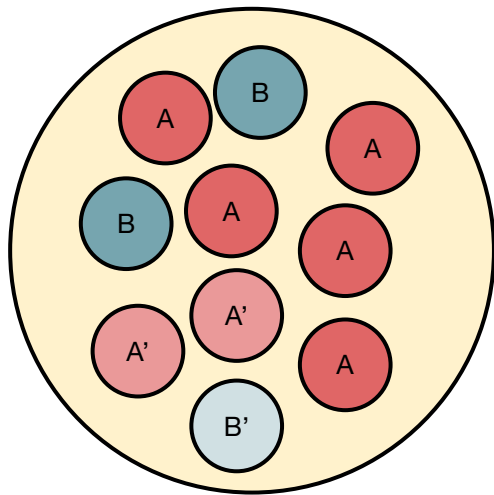
RQ1

Data Distribution Issues



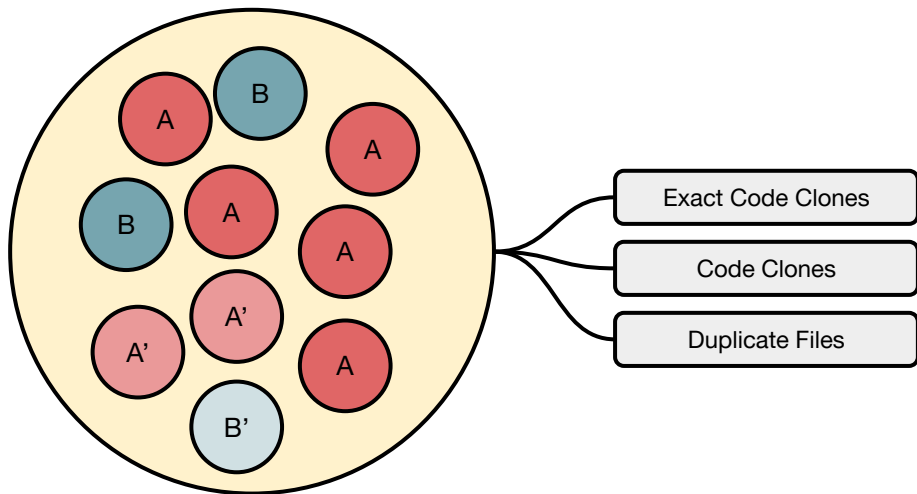
RQ1

(Near) Duplicated Instances



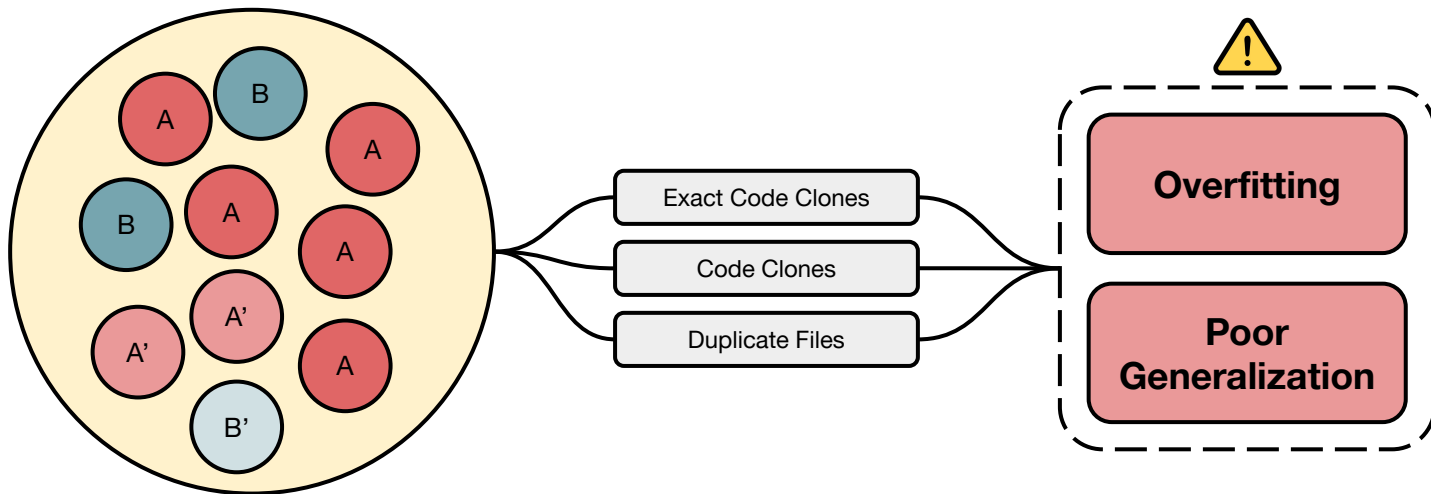
RQ1

(Near) Duplicated Instances



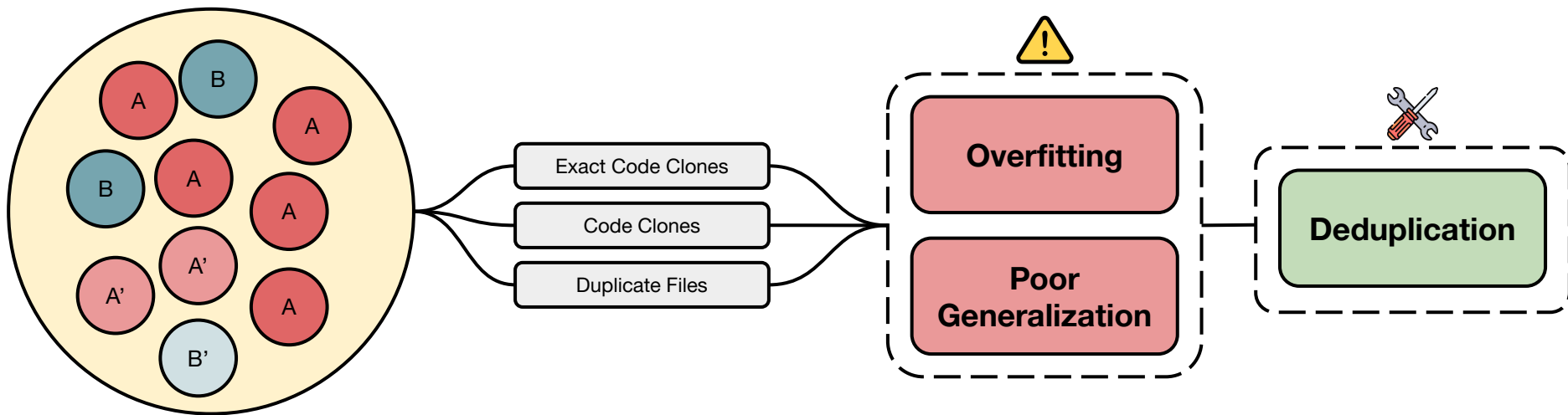
RQ1

(Near) Duplicated Instances



RQ1

(Near) Duplicated Instances



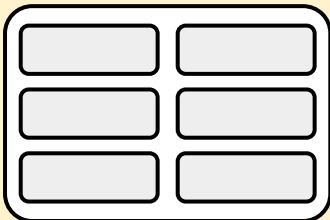
Natural Language

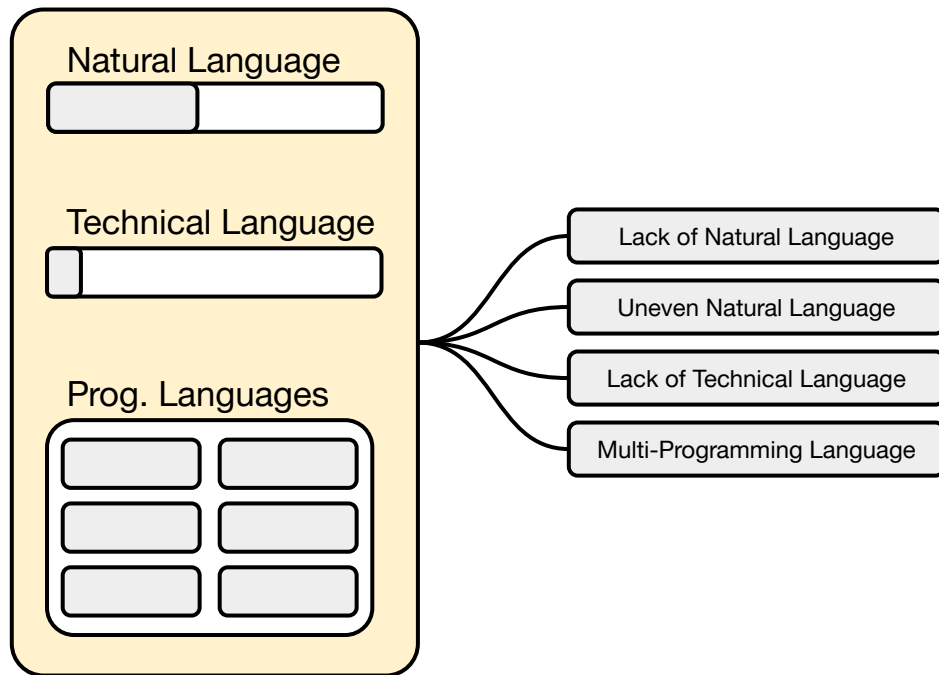


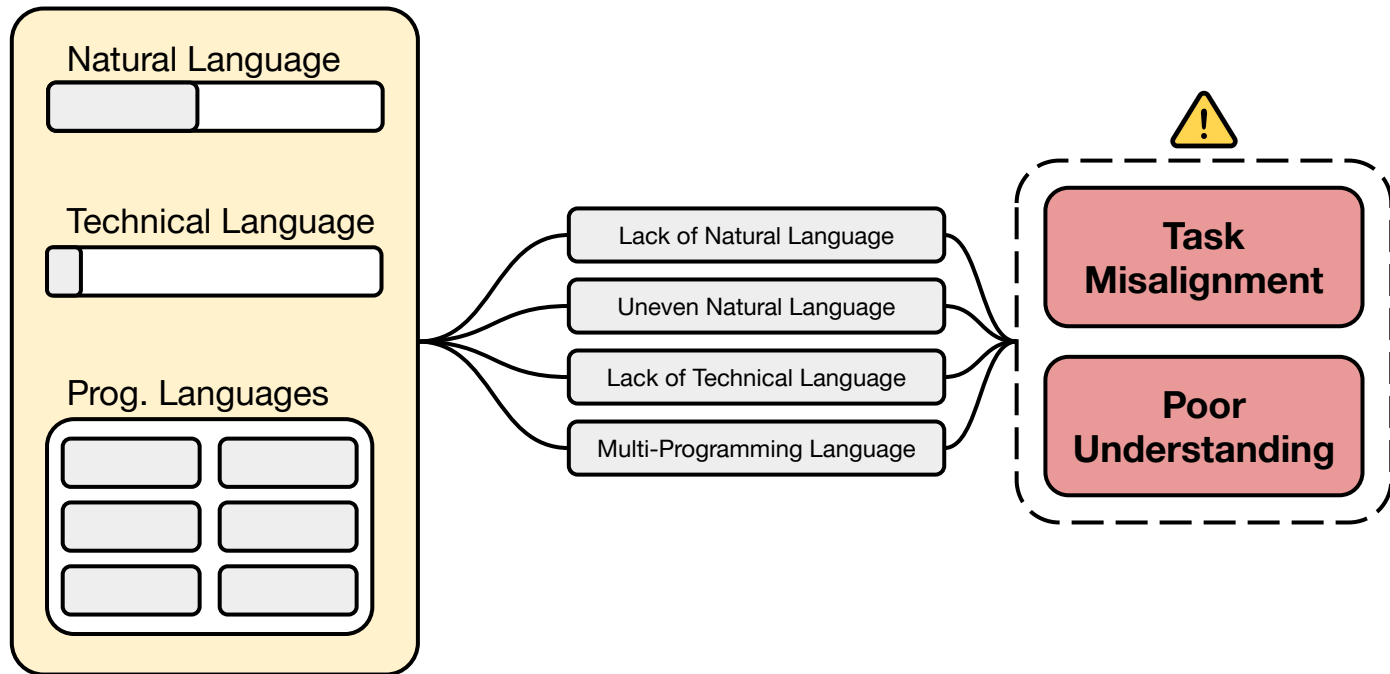
Technical Language

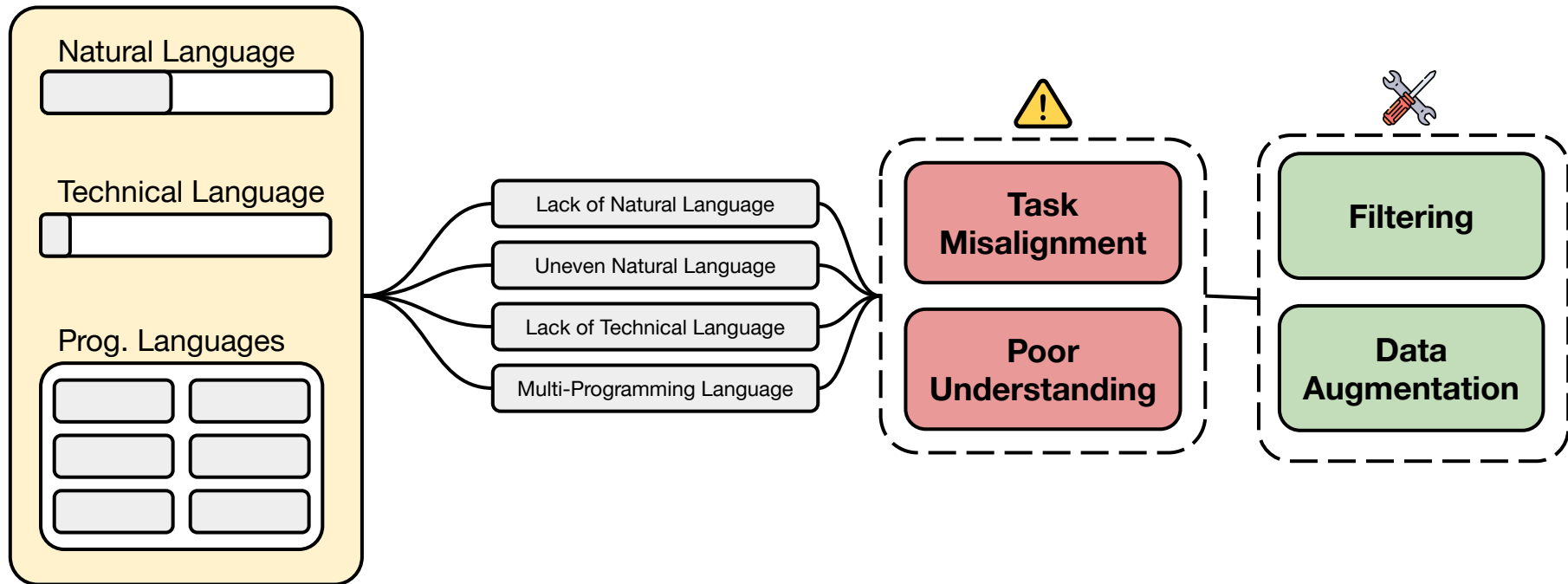


Prog. Languages



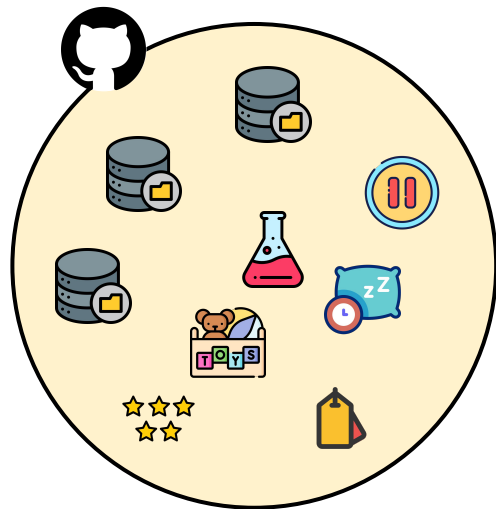






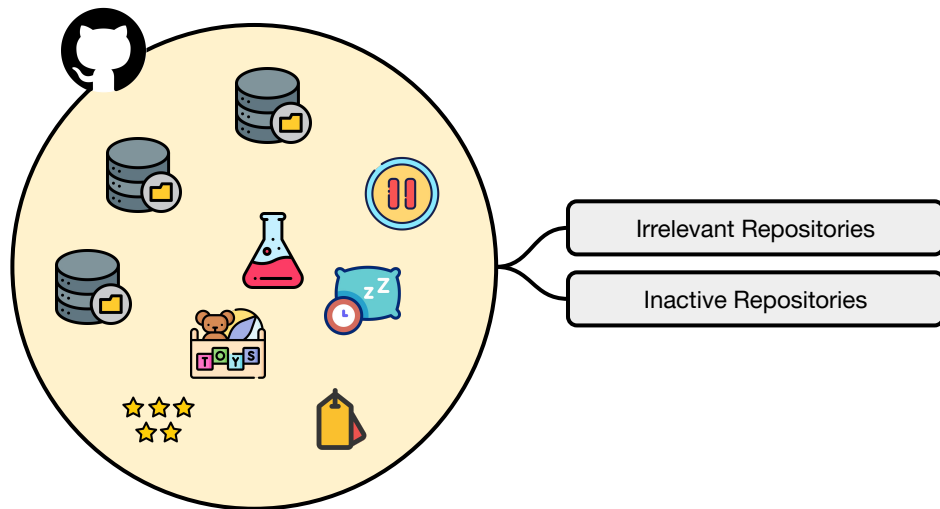
RQ1

Inadequate Source Repositories



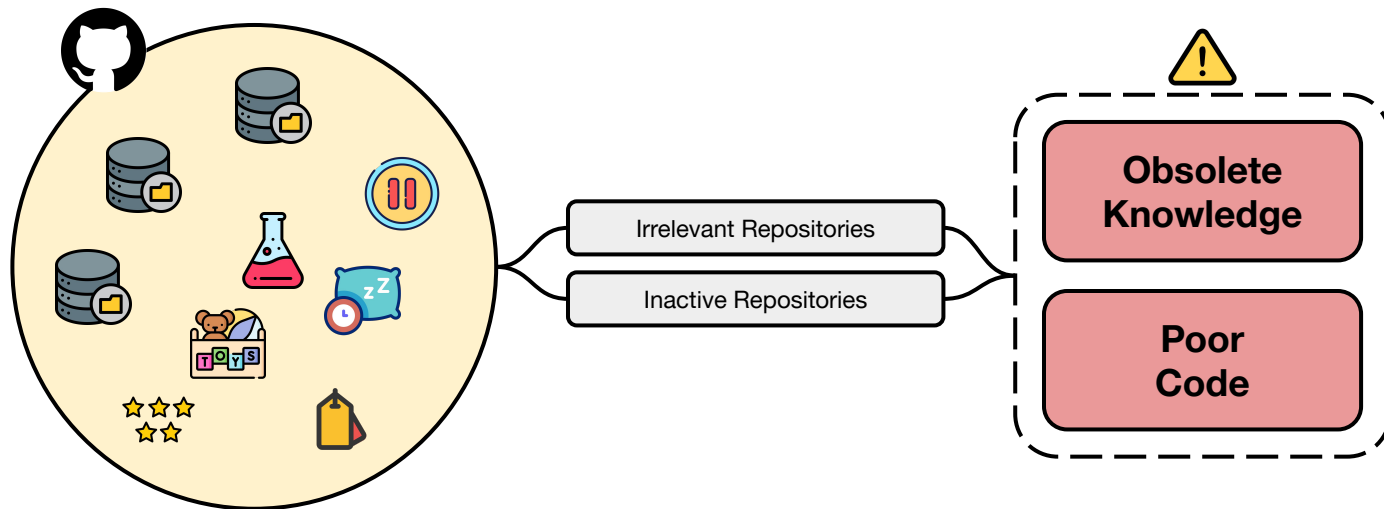
RQ1

Inadequate Source Repositories



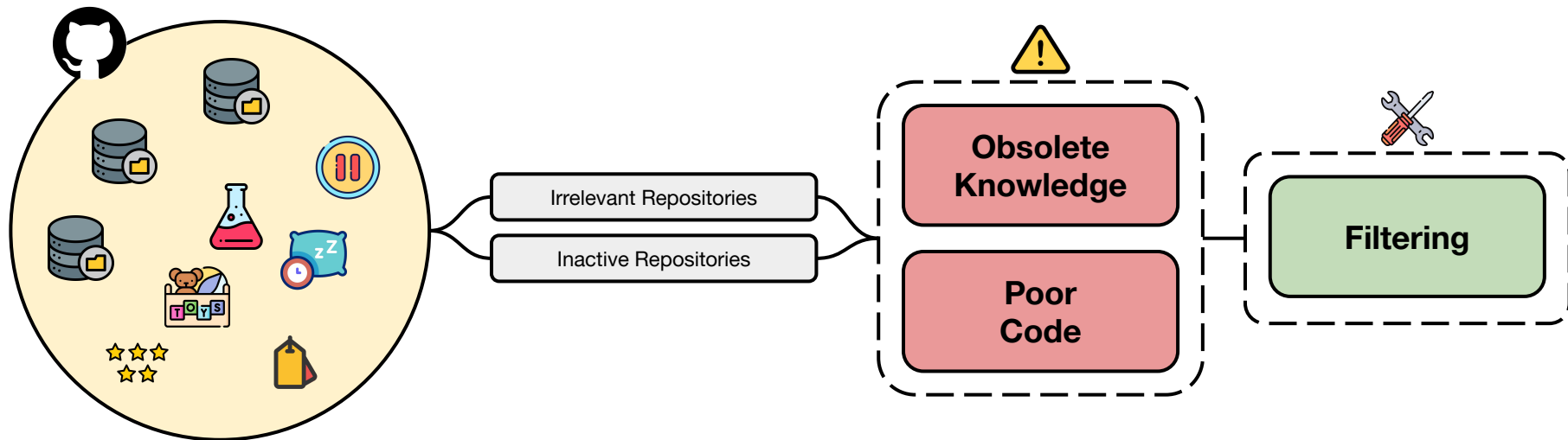
RQ1

Inadequate Source Repositories



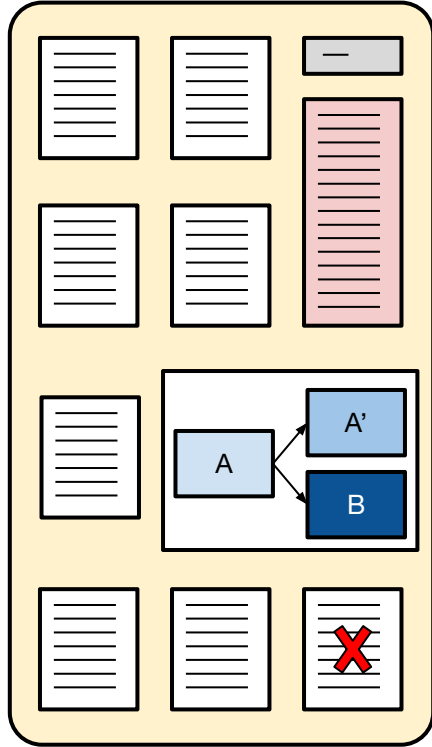
RQ1

Inadequate Source Repositories



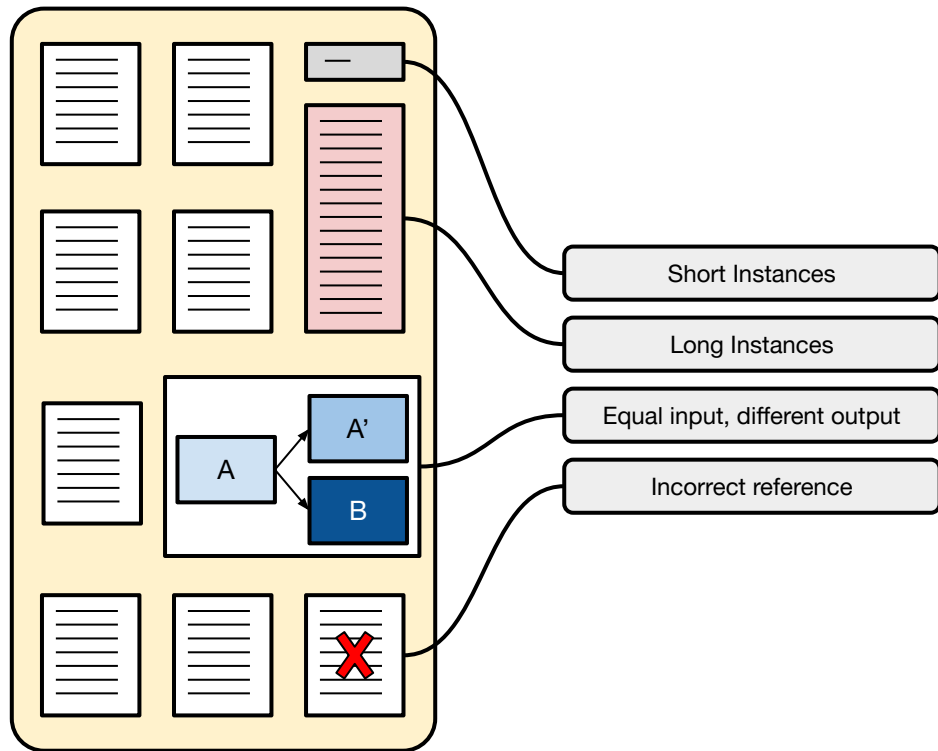
RQ1

Misleading Instances



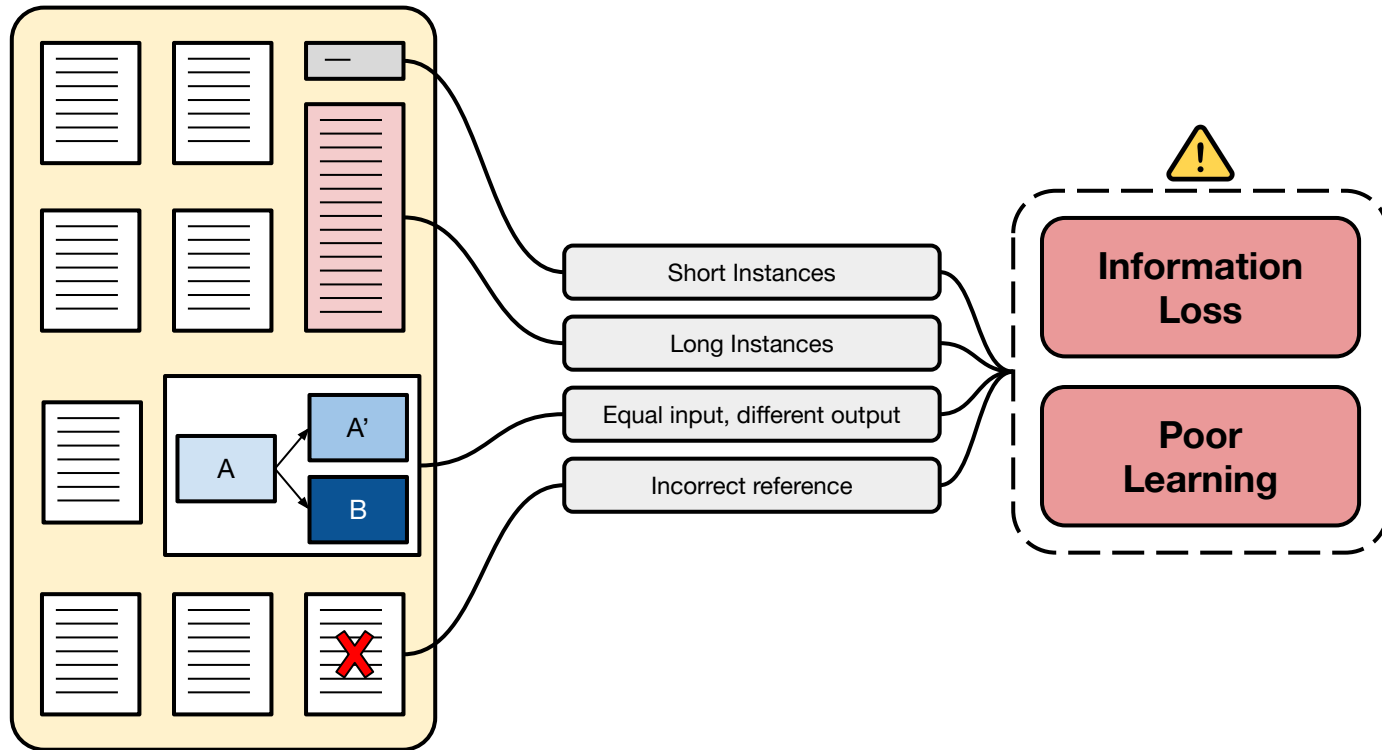
RQ1

Misleading Instances



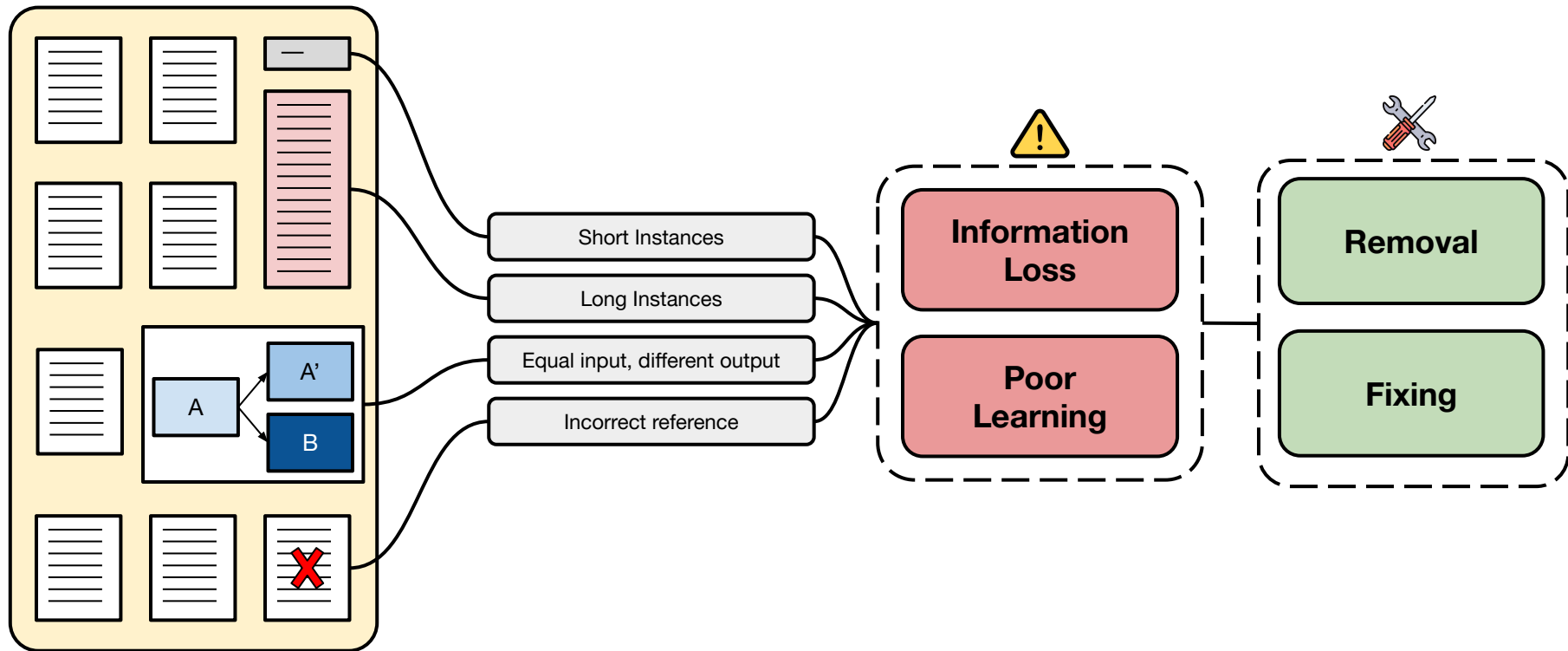
RQ1

Misleading Instances



RQ1

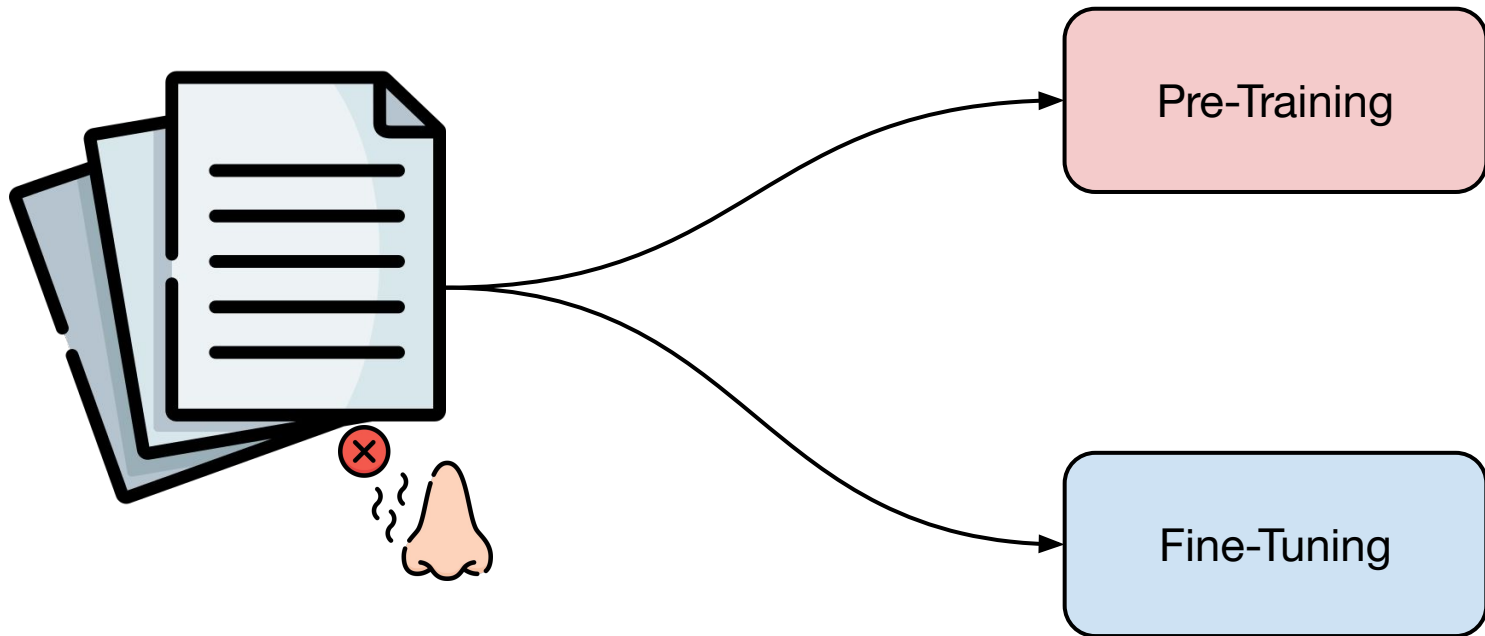
Misleading Instances



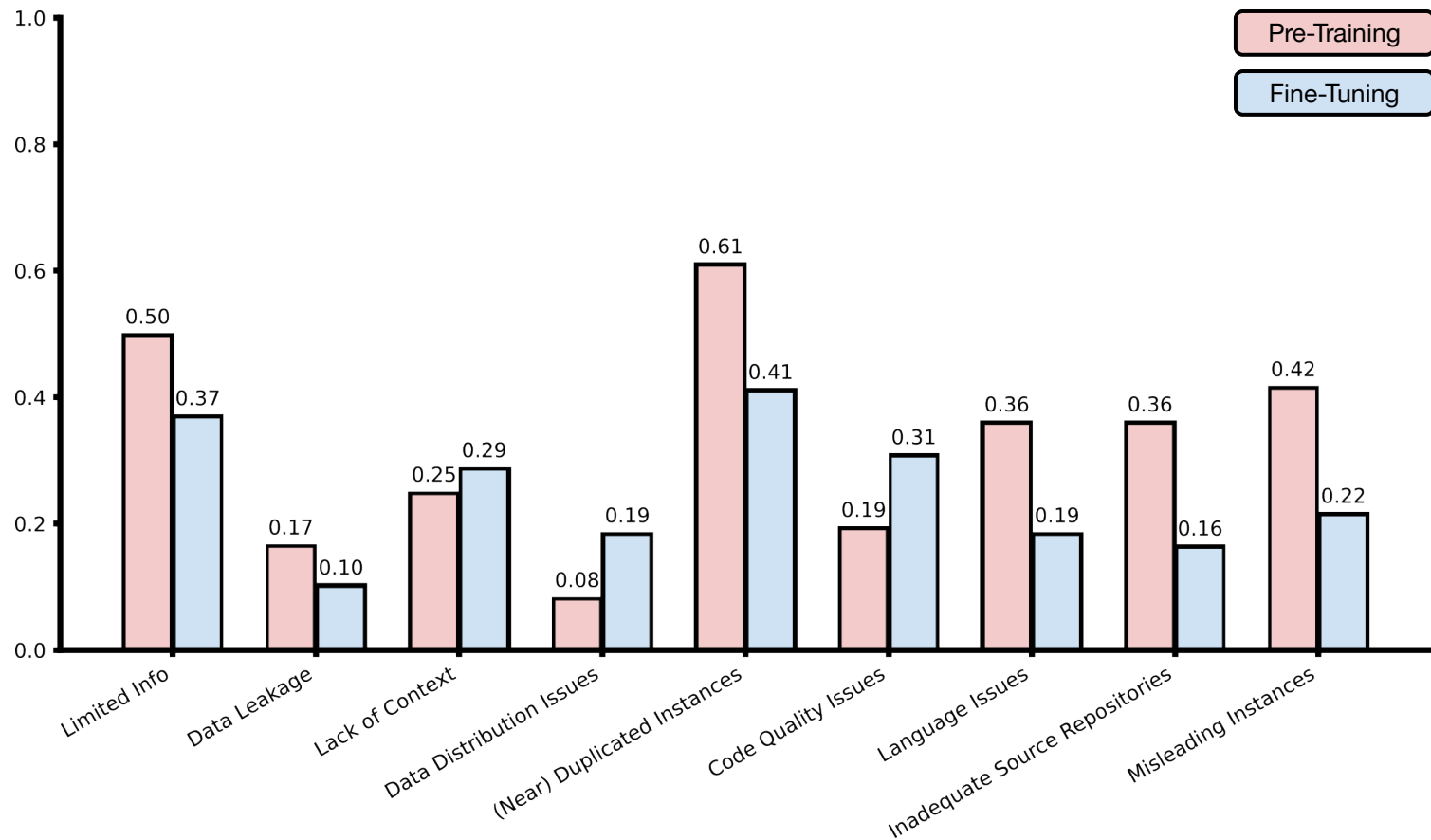
How **frequent** is the use
of strategies to **remove**
data smells between
pre-training and
fine-tuning



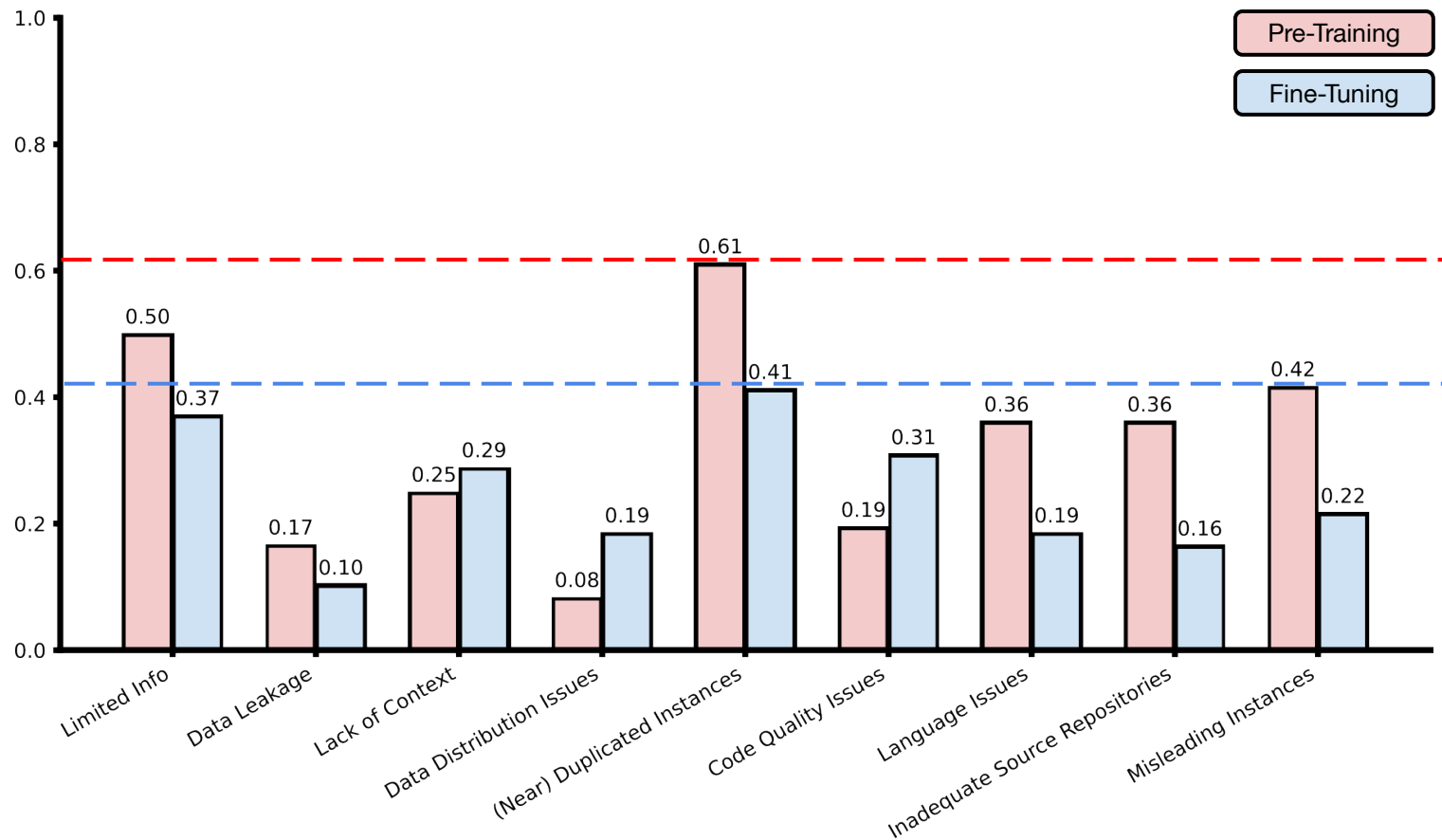
DESIGN



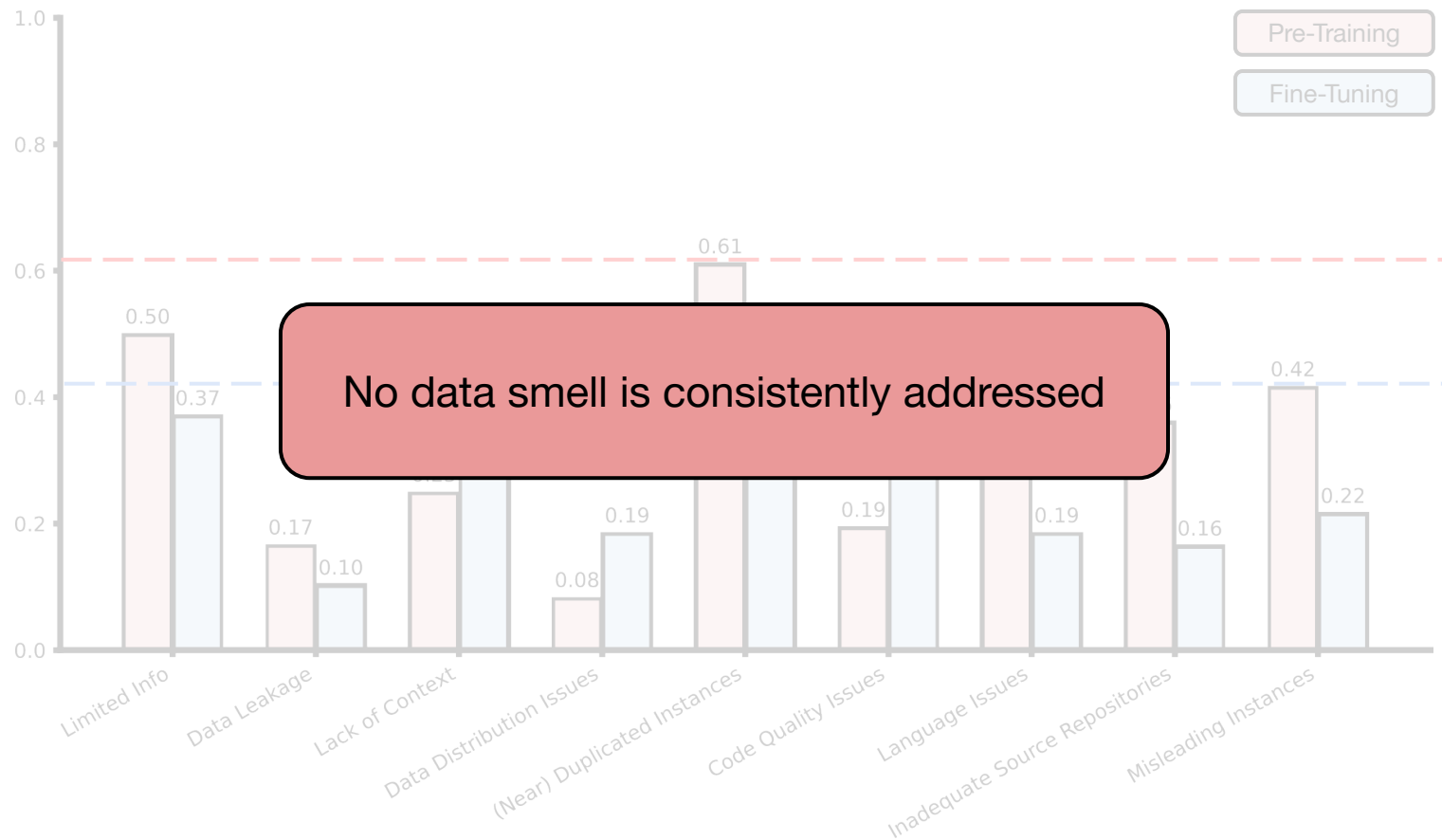
RQ2



RQ2



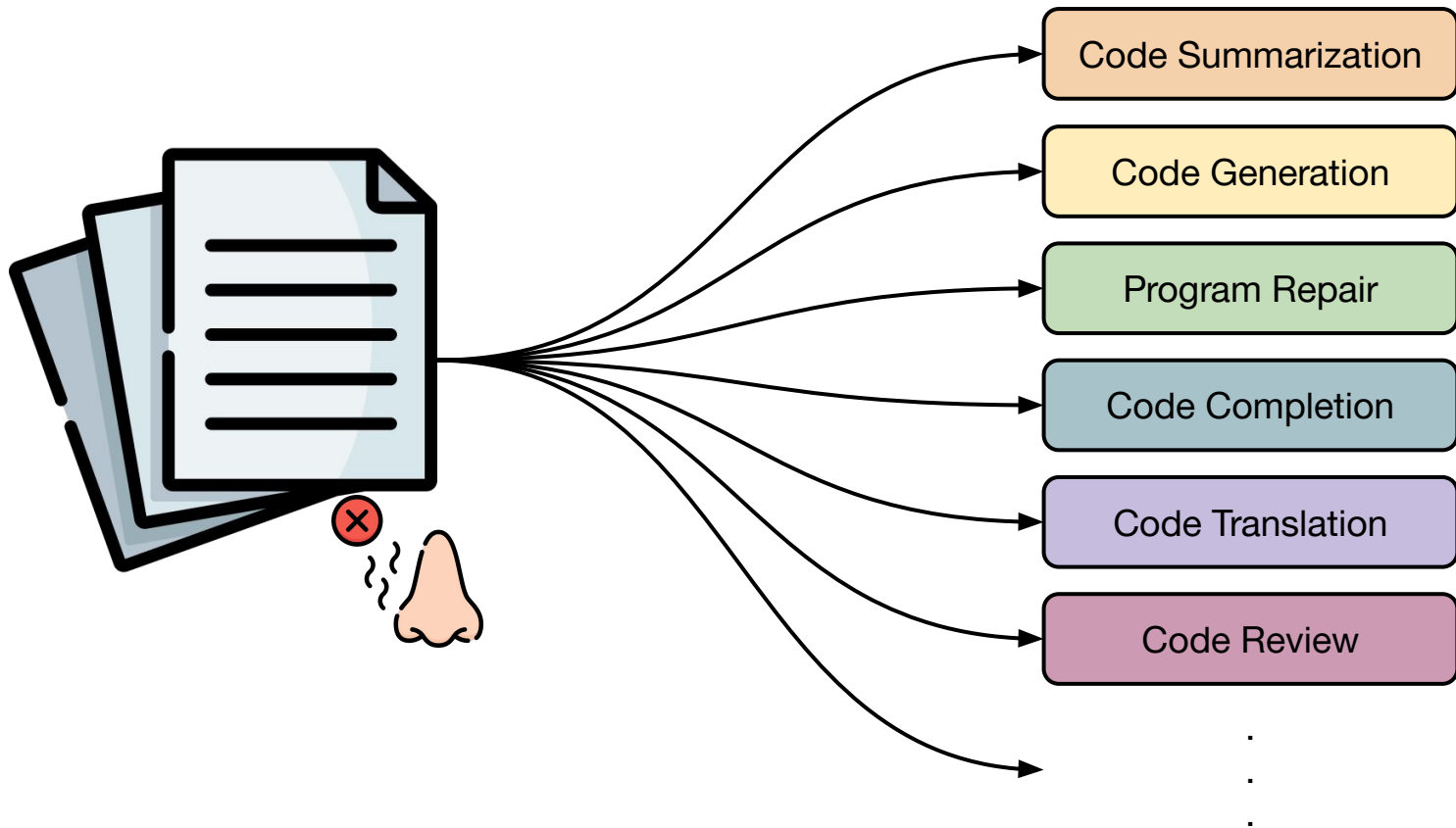
RQ2



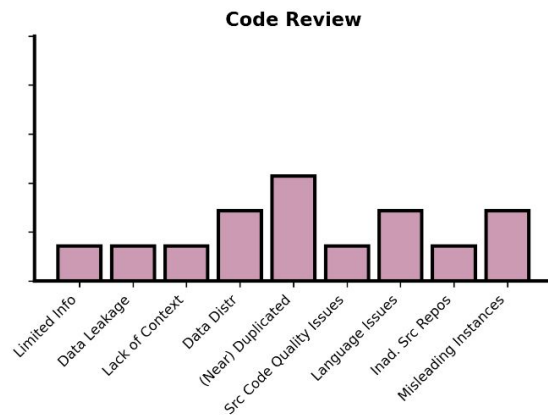
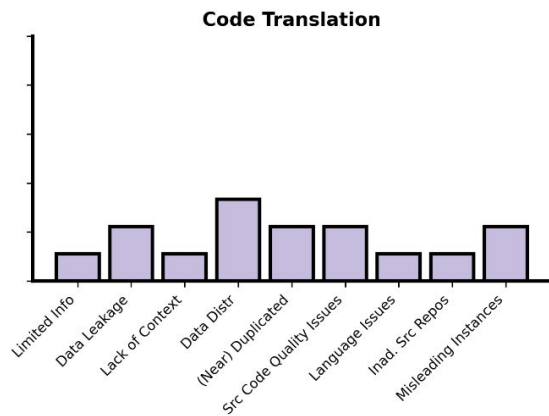
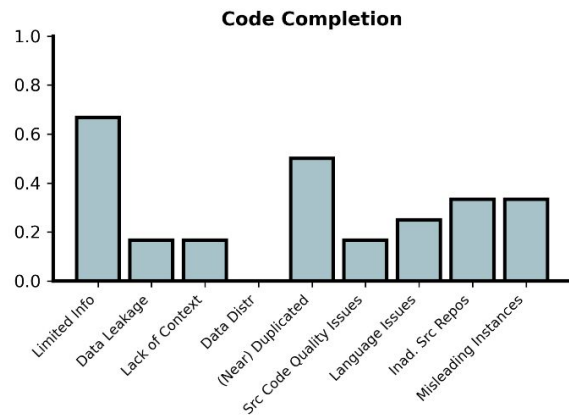
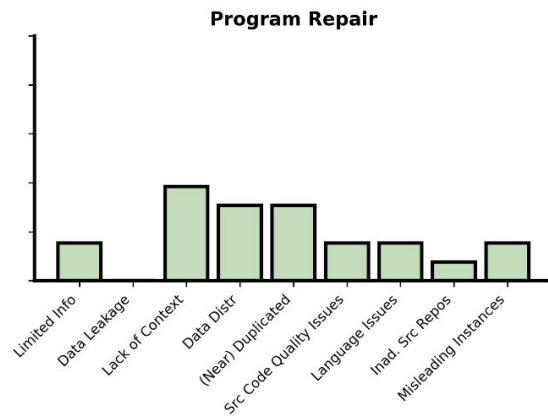
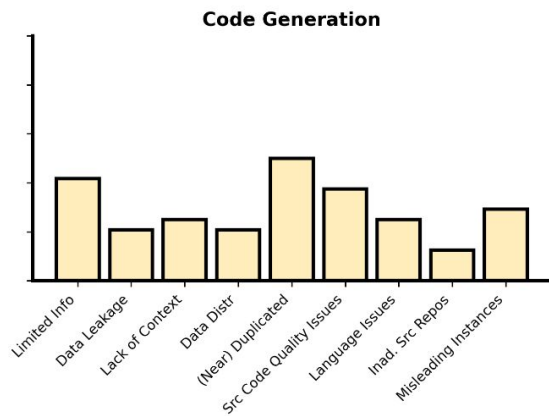
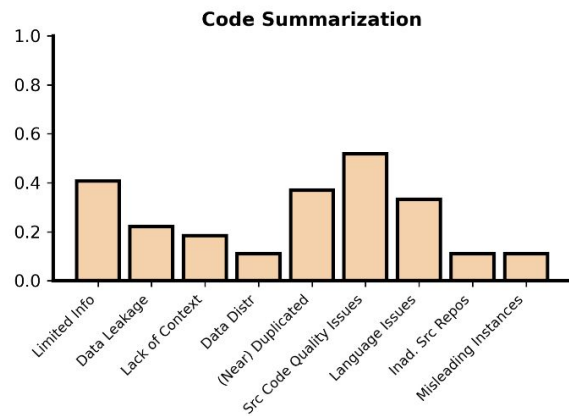
How **frequent** is the use
of strategies **to remove**
data smells among
different **coding tasks**



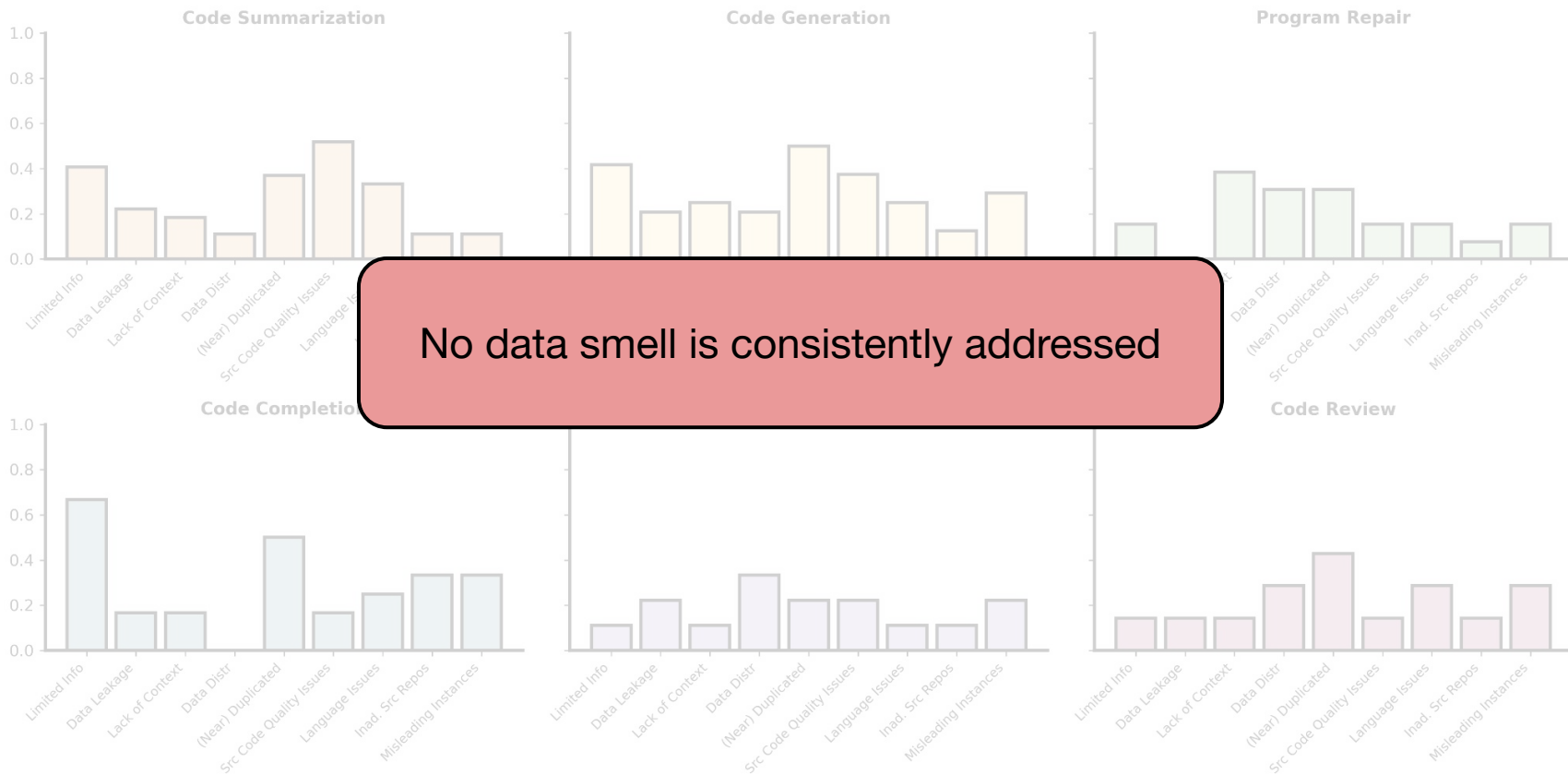
DESIGN



RQ3



RQ3



OUTCOMES

Discrepancy:

A discrepancy is found between the data smells addressed during **pre-training** and **fine-tuning** stages.

It is evidenced that many **code-related tasks DO not** take into account many data smells.

Shallowness:

Many cleaning techniques are **shallow** and, consequently, no precise.

Beyond Effectiveness:

It is needed to measure the impact of data smells also on **other aspects** (*i.e.*, security, privacy, robustness).

Quality Improvement:

To improve data quality, a canonical approach is to **remove low-quality instances**.

On the other hand, **augment instances missing features** can be crucial.

OUTCOMES

Discrepancy:

A discrepancy is found between the data smells addressed during **pre-training** and **fine-tuning** stages.

It is evidenced that many **code-related tasks DO not** take into account many data smells.

Shallowness:

Many cleaning techniques are **shallow** and, consequently, no precise.

Beyond Effectiveness:

It is needed to measure the impact of data smells also on **other aspects** (*i.e.*, security, privacy, robustness).

Quality Improvement:

To improve data quality, a canonical approach is to **remove low-quality instances**.

On the other hand, **augment instances missing features** can be crucial.

OUTCOMES

Discrepancy:

A discrepancy is found between the data smells addressed during **pre-training** and **fine-tuning** stages.

It is evidenced that many **code-related tasks DO not** take into account many data smells.

Shallowness:

Many cleaning techniques are **shallow** and, consequently, no precise.

Beyond Effectiveness:

It is needed to measure the impact of data smells also on **other aspects** (*i.e.*, security, privacy, robustness).

Quality Improvement:

To improve data quality, a canonical approach is to **remove low-quality instances**.

On the other hand, **augment instances missing features** can be crucial.

OUTCOMES

Discrepancy:

A discrepancy is found between the data smells addressed during **pre-training** and **fine-tuning** stages.

It is evidenced that many **code-related tasks DO not** take into account many data smells.

Shallowness:

Many cleaning techniques are **shallow** and, consequently, no precise.

Beyond Effectiveness:

It is needed to measure the impact of data smells also on **other aspects** (*i.e.*, security, privacy, robustness).

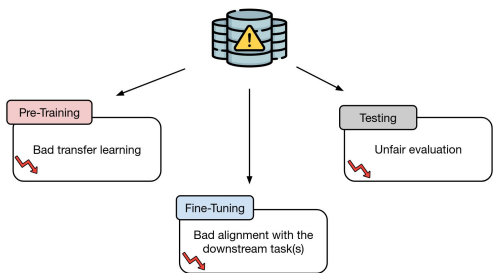
Quality Improvement:

To improve data quality, a canonical approach is to **remove low-quality instances**.

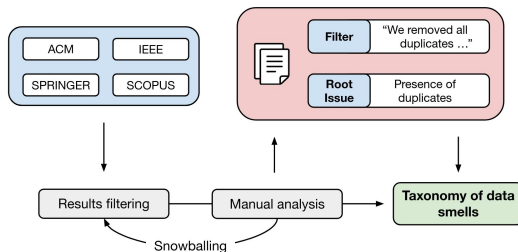
On the other hand, **augment instances missing features** can be crucial.

SUMMARY

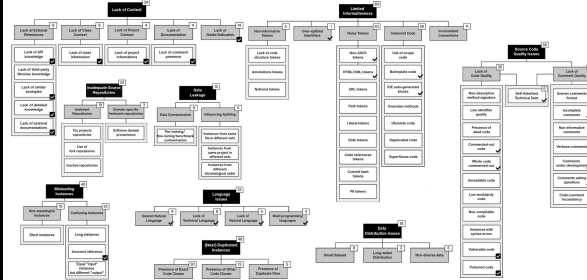
DATA IMPORTANCE



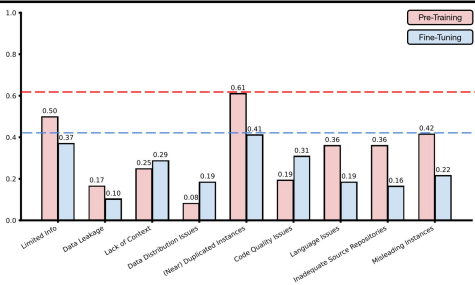
DESIGN



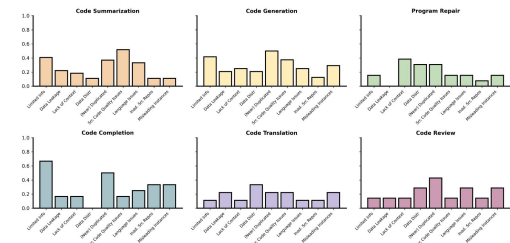
RQ1



RQ2



RQ3



OUTCOMES

Discrepancy:

A discrepancy is found between the data smells addressed during **pre-training** and **fine-tuning** stages.

Shallowness:

Many cleaning techniques are **shallow** and, consequently, not precise.

Quality Improvement:

To improve data quality,
a canonical approach is
to **remove low-quality
instances.**

On the other hand, **augment instances missing features** can be crucial.

Beyond Effectiveness:

It is needed to measure the impact of data smells also on **other aspects** (i.e., security, privacy, robustness).