

Quantifying the benefits of code hints for refactoring deprecated Java APIs

Cristina David
University of Bristol
Bristol, UK

Daniel Kroening
Amazon Inc.
Seattle, US

Pascal Kesseli
Meta
Zurich, Switzerland

Hanliang Zhang
University of Bristol
Bristol, UK

Abstract

When done manually by engineers at Amazon and other companies, refactoring legacy code in order to eliminate uses of deprecated APIs is an error-prone and time-consuming process. In this paper, we investigate to which degree refactorings for deprecated Java APIs can be automated, and quantify the benefit of Javadoc code hints for this task. To this end, we build a symbolic and a neural engine for the automatic refactoring of deprecated APIs. The former is based on type-directed and component-based program synthesis, whereas the latter uses LLMs. We applied our engines to refactor the deprecated methods in the Oracle JDK 15. Our experiments show that code hints are enabling for the automation of this task: even the worst engine correctly refactors 71% of the tasks with code hints, which drops to at best 14% on tasks without. Adding more code hints to Javadoc can hence boost the refactoring of code that uses deprecated APIs.

CCS Concepts

• Software and its engineering → Empirical software validation; Source code generation; • Computing methodologies → Machine learning.

Keywords

Program Refactoring, Program Synthesis, LLMs

ACM Reference Format:

Cristina David, Pascal Kesseli, Daniel Kroening, and Hanliang Zhang. 2025. Quantifying the benefits of code hints for refactoring deprecated Java APIs. In *33rd ACM International Conference on the Foundations of Software Engineering (FSE Companion '25)*, June 23–28, 2025, Trondheim, Norway. ACM, New York, NY, USA, 12 pages. <https://doi.org/10.1145/3696630.3728567>

1 Introduction

As programming languages evolve, industry engineers take on the critical task of migrating code to newer versions, an important aspect of code modernisation. This challenge spans a wide range of companies, including Amazon. Such migrations ensure access to new features, improved performance, and enhanced security while maintaining compatibility with modern tools and libraries. Java, for

example, has frequent JDK releases, each introducing improvements and updates. However, migrating to newer JDK versions presents challenges, particularly due to deprecated APIs. Deprecation, intended to enhance security, performance, and maintainability, requires engineers to update existing code by replacing deprecated APIs with modern, more efficient alternatives, adding complexity to the migration process.

The transformation of existing code such that it doesn't use deprecated APIs is not always straightforward, as illustrated in Figure 1. In the example, we make use of the `getHours` method of the `Date` class, which is deprecated. In this situation, in order to replace the use of the deprecated method, we must first obtain a `Calendar` object. However, we can't use the `Calendar` constructor as it is protected, and we must instead call `getInstance`. Furthermore, in order to be able to use this `Calendar` object for our purpose, we must first set its time using the existing date. We do this by calling `setTime` with `date` as argument. Finally, we can retrieve the hour by calling `calendar.get(Calendar.HOUR_OF_DAY)`.

```
void main(String[] args) {  
    // Deprecated:  
    int hour = date.getHours();  
  
    // Should have been:  
    final Calendar calendar = Calendar.getInstance();  
    calendar.setTime(date);  
    int hour = calendar.get(Calendar.HOUR_OF_DAY);  
}
```

Figure 1: Deprecated method example.

Refactoring code that relies on deprecated APIs presents additional challenges. For instance, the code to be refactored might be using abstract classes and abstract methods, and it may not be obvious how to subclass from the code to be refactored (e.g. `engineGetParameter` in `java.security.SignatureSpi`). It may also call methods that, while not abstract, must be overridden by subclasses (e.g., the `layout` method in `java.awt.Component`, which has an empty body). Furthermore, the code might invoke native methods implemented in other programming languages, such as C or C++ (e.g., `weakCompareAndSet` in `java.util.concurrent.atomic.AtomicReference`). Understanding the behavior of such code requires knowledge of how to subclass abstract classes, override methods appropriately, and understand the functionality of native code written in other languages.



This work is licensed under a Creative Commons Attribution 4.0 International License. *FSE Companion '25, Trondheim, Norway*
© 2025 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-1276-0/2025/06
<https://doi.org/10.1145/3696630.3728567>

```

/*
 * @deprecated As of JDK version 1.1,
 * replaced by {@code Calendar.get
 *             (Calendar.HOUR_OF_DAY)}.
 */

```

Figure 2: Code hints for the running example.

Manual refactoring of deprecated APIs is a time-consuming and error-prone process. *Therefore, in this paper, we investigate the automatic generation of refactorings for deprecated APIs.* While our primary emphasis is on refactoring deprecated methods, the same techniques can be extended to handle deprecated fields and classes.

When deprecating a field, method, or class, the `@Deprecated` Javadoc tag is used in the comment section to inform the developer of the reason for deprecation and, sometimes, what can be used in its place. We call such a suggestion a *code hint*. These hints are helpful for a human performing the refactoring. While they don't offer a complete solution, they serve as valuable guidance, making the process more intuitive.

For illustration, let's look at our running example in Figure 1. The source code for the `getHours` method in class `Date` is accompanied by the comment in Figure 2, where the `@code` tag suggests replacing the deprecated `getHours` by `Calendar.get(Calendar.HOUR_OF_DAY)`. Using this code hint is not straightforward. Although it may seem as if method `get` is static, allowing an immediate call, it is actually an instance method, and requires an object of class `Calendar`. However, no such object is available in the original code, meaning that it must be created by the refactored code. Consequently, we must find the necessary instructions that consume existing objects and create a `Calendar` object. Besides generating the required objects, we also need to set their fields. For instance, in Figure 1, we must call `calendar.setTime(date)` to set the calendar's date based on the existing date object.

Given that code hints are useful to humans, *in this paper, we are also interested in investigating the benefits of Javadoc code hints in the automation of the refactoring.*

In order to investigate the benefits of code hints when automating the deprecated refactoring, we build two code generation engines, namely a symbolic and a neural one. For each, we selected approaches with proven success in program synthesis, and that could also incorporate code hints. In particular, for the symbolic engine, we make use of component-based synthesis [12, 16, 26] and type-directed synthesis [30, 42, 57]. Essentially, we use types and code hints to populate a component library (i.e., a library of instructions), such that these components are then weaved together to generate the desired program. Given the recent success of Large Language Models (LLMs) for code generation [8, 24, 44, 56, 58–60, 62], the neural approach generates candidate refactorings by iteratively querying an LLM – we used Claude 2.1 and Claude 3 [1].

Our experiments show that code hints are important for the performance of both the symbolic and the neural engines: when code hints are given, both engines perform very well, making full automation of deprecated APIs refactorings feasible. Without code hints, the refactoring becomes much harder for both engines, such

that the vast majority of benchmarks without code hints are failing. Thus, our conclusion is that, in order to facilitate the automation of the refactoring of client code that uses deprecated APIs, code hints should be added to all deprecated methods that have a replacement.

Comparing the symbolic and the neural approaches, code hints help the symbolic approach to efficiently prune the solution space, resulting in a better performance than the neural engine at a lower computational cost. We believe this is important to note, especially as LLMs are often seen as a panacea for all tasks, including code generation. This work shows that symbolic methods can still be effective in specialised settings, where efficient pruning of the solution space is possible.

Contributions.

- We propose a symbolic and a neural refactoring technique for deprecated APIs. The former makes use of type-directed and component-based synthesis, whereas the latter is LLM based. In order to check the correctness of the refactorings, we design a symbolic equivalence check, which takes into consideration both the state of the stack and the heap.
- We implement our techniques and use them to refactor deprecated methods from the Oracle JDK 15 Deprecated API documentation [41].
- We investigate the benefits of code hints for both the symbolic and neural approach. Our results show that code hints are important for the performance of both the symbolic and the neural engines: adding more code hints to Javadoc can substantially help automate the refactoring of deprecated APIs.

2 Our approach

In this section, we describe the two code generation engines. For both approaches, we make use of a CounterExample Guided Inductive Synthesis (CEGIS) [48] architecture, where we iteratively attempt to improve a candidate refactoring until it behaves indistinguishably from the original code, i.e., the original and the refactored blocks of code are observationally equivalent. In the rest of the paper, we will use the predicate $\text{equivalent}(P_1(\vec{i}), P_2(\vec{i}))$ to denote that code P_1 is observationally equivalent to P_2 for inputs $\vec{i}$, meaning that the two programs can't be distinguished by their behaviour on inputs $\vec{i}$. In general, we refer to the original code as P_1 and the refactored code as P_2 . We will fully define what *equivalent* actually means in Section 2.3.

In each iteration of the synthesis process, there are two phases, a synthesis phase and a verification phase. The synthesis phase generates a *candidate refactoring* that is equivalent to the original code on a finite set of inputs. Then, the verification phase tries to find a new *counterexample input* that distinguishes between the current candidate refactoring and the original code. If it manages, this input is added to the set of finite inputs used by the next synthesis phase, and if it fails, then the current candidate is indeed a solution refactoring.

2.1 Synthesis phase

In this phase, we have a finite set of input examples $\{\vec{i}_1 \dots \vec{i}_n\}$ and we attempt to find a candidate refactoring that is observationally equivalent to the original code for the given inputs.

Symbolic engine. We construct the following Synthesise method, which takes the refactored code P_2 as input.

```
Synthesise ( $P_2$ ) {
  if (equivalent( $P_1(\vec{i}_1)$ ,  $P_2(\vec{i}_1)$ ) && ...
    && equivalent( $P_1(\vec{i}_n)$ ,  $P_2(\vec{i}_n)$ ))
    assert(false);
}
```

This method consists of only one conditional saying that if P_1 and P_2 are equivalent on all given inputs $\vec{i}_1, \dots, \vec{i}_n$, then `assert(false)` is reached. Given that this assertion always fails if reached, it means that the method is unsafe when P_1 and P_2 are equivalent on the given inputs. Thus, we can reduce the synthesis problem to the problem of checking the safety of Synthesise. If a safety checker manages to find an input P_2 for which the assertion fails (meaning that Synthesise is unsafe), then this P_2 must be equivalent to P_1 on the given inputs. This P_2 is the *refactoring candidate* that the synthesis phase is supposed to generate.

We use an existing fuzz testing platform for Java, JQF [43], as the safety checker. JQF is designed to handle structured inputs, where inputs of type T are generated by a backing `Generator<T>`. JQF provides a library of generators for basic types such as primitive values. We implement custom JQF generators based on a library of instructions built as explained in Section 3, enabling JQF to construct P_2 by weaving together instructions from the library.

If the safety checker fails to find a P_2 for which the assertion fails, then the overall synthesis technique fails to generate a solution refactoring. There could be two reasons for this situation, which we cannot differentiate between. Firstly, there may indeed be no P_2 that can be constructed with the available components in the component library such that it is equivalent to P_1 on the given inputs. Secondly, this could be caused by JQF's unsoundness. Given that fuzzers rely on testing, they may fail to find inputs that trigger unsafe behaviours even when such inputs exist. Consequently, we cannot guarantee that we will always find a refactoring whenever one exists. We will provide more details about our decision to use fuzzing as the safety checker in Section 4.

Neural engine. For each deprecated method, we query the LLM by constructing a prompt as given in Figure 3. The method definition and the JavaDoc comment (if present) are provided as context. Furthermore, we provide a code snippet of the use of the deprecated method that is to be refactored (e.g., `this.minimumSize();`). Finally, we attach a list of formatting constraints.

The code returned by the LLM is verified against the original code. If the verification fails, we attach the set of counterexamples generated by the fuzzing engine to subsequent prompts. We apply object serialisation to obtain structured-views of the input and output states. In addition to this, we followed the prompting guidelines provided by Anthropic Claude¹ (e.g. using XML tags). Do note that, as opposed to the symbolic approach, the LLM doesn't

¹<https://docs.anthropic.com/en/docs/prompt-engineering>

Initial Prompt
User: <i># Context</i> The method <code>\$(METHOD_NAME)</code> of the class <code>\$(CLASS_NAME)</code> is deprecated. Below is the method definition: <code>\$(METHOD_DEFINITION)</code> Here are its Javadoc comments that may contain a <code>@deprecated</code> tag explaining why the item has been deprecated and suggesting what to use instead. <code>\$(JAVADOC_COMMENT)</code> However, I used this method call in my code base, the code snippet is given below: <code>\$(CODE_SNIPPET)</code> <i># Instruction</i> Help me refactor this code snippet so that it doesn't use the deprecated method. Do not simply inline the method body, use APIs suggested by the Javadoc comments if there are any. <i># Constraints</i> Take the following constraints into consideration: <code>\$(FORMATTING_CONSTRAINTS)</code>
Subsequent Prompt
Additionally, here is a set of input/output examples that you should respect, <code>\$(EXAMPLES)</code> Assistant:

Figure 3: LLM Prompt Template.

give any guarantees that the counterexamples were actually taken into consideration.

2.2 Verification phase

This is exactly the same for the two approaches. We are provided with a candidate refactoring P_2 and we must check whether there exists any input $\vec{i}$ for which the original code and the candidate refactoring are not observationally equivalent. To do this, we build the following Verify method, which given some input $\vec{i}$, asserts that the two programs are equivalent for $\vec{i}$.

```
Verify( $\vec{i}$ ) { assert(equivalent( $P_1(\vec{i})$ ,  $P_2(\vec{i})$ )); }
```

In other words, answering the question posed by the verification phase is reduced to checking the safety of this method: if Verify is safe (i.e., the assertion is not violated) for any input $\vec{i}$, then there is no input that can distinguish between the original code and the candidate refactoring (this is indeed a sound refactoring). However, if Verify is not safe, then we want to be able to obtain a *counterexample input* $\vec{i}_{cex}$ for which the assertion fails. This counterexample will be provided back to the synthesis phase and used to refine the current candidate refactoring. Again, we check the safety of Verify with JQF.

In this section, we hid the complexity of the equivalence check inside the *equivalent* predicate. This is far from trivial as it requires handling of notions such as aliasing, loaded classes, static fields etc. This will be described in detail next.

2.3 Checking program equivalence

A core part of the synthesis procedure is the *equivalent* $(P_1(\vec{i}), P_2(\vec{i}))$ predicate, which checks that the original code P_1 and a candidate refactoring P_2 are equivalent for a given input $\vec{i}$. In this section, we provide details on how we check this equivalence.

The state of a Java program is modelled by the current program stack (consisting of method-specific values and references to objects in the heap) as well as its heap (consisting of instance variables and static field values). Static field values are stored with their respective classes, which in turn are loaded by class loaders. Our equivalence check must take all these into consideration. One of the main challenges is aliasing, which we will discuss later in the section.

We start by introducing some notation:

- $loadedClasses(P)$ returns the set of classes loaded by the class loader in which P is executed. Note that Java allows to load the same class in different class loaders, which creates independent copies of its static fields.
- $liveVars(P)$ provides the set of variables that are live at the end of P .
- $staticFields(C)$ returns the set of static fields of class C .
- $exc(P, \vec{i})$ returns the exception thrown by P 's execution on input $\vec{i}$.

EXAMPLE 1. For our running example in Figure 1, P_1 and P_2 are represented by the following lines of code. Note that both P_1 and P_2 take variable *date* as input.

```
// P1
int hour=date.getHours();

// P2
final Calendar calendar=Calendar.getInstance();
calendar.setTime(date);
int hour=calendar.get(Calendar.HOUR_OF_DAY);
```

Then, we have:

```
liveVars(P1) = {date, hour}
liveVars(P2) = {calendar, date, hour}
loadedClasses(P1) = {Date}
loadedClasses(P2) = {Calendar, Date}
staticFields(Date) = {}
staticFields(Calendar) = {DATE, YEAR, ...}
exc(P1, _) = exc(P2, _) = {}
```

In order to extract the (last) object assigned to a variable v by the execution of P on a specific input $\vec{i}$, we will use the notation $E[P(\vec{i})](v)$. Essentially, if we consider the trace generated by executing $P(\vec{i})$, then $E[P(\vec{i})]$ maps each variable defined in P to the last object assigned to it by this trace. Next, we define the notion of equivalence with respect to a concrete input $\vec{i}$ (this definition is incomplete and we will build on it in the rest of the section). We overload equality to work over sets (for live variables and loaded classes).

DEFINITION 1 (PROGRAM EQUIVALENCE WITH RESPECT TO A CONCRETE INPUT $\vec{i}$ [partial]). Given two code blocks P_1 and P_2 and concrete input $\vec{i}$, we say that P_1 and P_2 are equivalent with respect

to $\vec{i}$, written as *equivalent* $(P_1(\vec{i}), P_2(\vec{i}))$ if and only if the following conditions hold:

- (1) $exc(P_1, \vec{i}) = \{e_1\} \wedge exc(P_2, \vec{i}) = \{e_2\} \wedge equals(e_1, e_2) \vee exc(P_1, \vec{i}) = exc(P_2, \vec{i}) = \emptyset$
- (2) $liveVar(P_1) = liveVar(P_2) \wedge \forall v \in liveVar(P_1). equals(E[P_1(\vec{i})](v), E[P_2(\vec{i})](v))$
- (3) $loadedClasses(P_1) = loadedClasses(P_2) \wedge \forall C \in loadedClasses(P_1). \forall f \in staticFields(C). equals(E[P_1(\vec{i})](f), E[P_2(\vec{i})](f))$

The above definition says that in order for P_1 and P_2 to be equivalent with respect to input $\vec{i}$, (1) either they both throw an exception and the two exceptions have the same type, or none does, (2) the set of variables live at the end of P_1 is the same as the set of variables live at the end of P_2 , and must be assigned equal objects by the executions of P_1 and P_2 on $\vec{i}$, respectively, and (3) the classes loaded by P_1 must be the same as the classes loaded by P_2 , and all the static fields in the classes loaded by both the class loaders of P_1 and P_2 must be assigned equal objects by the two executions, respectively. In our implementation, if either the deprecated or the refactored code didn't load a class, we load it for it, and check that the initial state of that class is the same for both blocks of code.

Equality refers to value equality, which we check by recursively following attribute chains until we reach primitive types. We generally do not consider existing equals methods unless (i) the method was not written by the user (i.e., JCL classes), (ii) the type inherits from `java.lang.Object`, (iii) the class implements `java.lang.Comparable` and (iv) the type declares an equals implementation. Examples of classes that satisfy these strict requirements are `java.lang.Integer` or `java.util.Date`. All other implementations of equals are considered unreliable and ignored.

EXAMPLE 2. When applying Definition 1 to P_1 and P_2 given in Example 1, the following conditions must hold (given that class *Date* has no static fields, the third condition is trivial):

- (1) $exc(P_1) = exc(P_2) = \emptyset$
- (2) $equals(E[P_1(\vec{i})](hour), E[P_2(\vec{i})](hour)) \wedge equals(E[P_1(\vec{i})](date), E[P_2(\vec{i})](date))$
- (3) $Date \in \{Date, Calendar\}$

The challenge of aliasing. When expressing the equivalence relation between P_1 and P_2 , we intentionally missed one important aspect, namely aliasing. To understand the problem let's look at the following example, which makes use of `java.awt.Container`, where a generic Abstract Window Toolkit (AWT) container object is a component that can contain other AWT components. Method `preferredSize` used by the original code below returns the preferred size of the calling container. Method `preferredSize` is deprecated, and, instead, the refactored version uses `getPreferredSize`.

EXAMPLE 3.

```
// P3:
Dimension dim1=container.preferredSize();
Dimension dim2=container.preferredSize();

// P4:
```

```
Dimension dim1=container.getPreferredSize();
Dimension dim2=dim1;
```

We note that, the original code above defines two variables *dim1* and *dim2*, each assigned an object of type *Dimension* returned by calling *container.preferredSize()*. Conversely, in the refactored code, *dim1* and *dim2* are aliases, i.e., they point to the same object of type *Dimension*.

The original and the refactored code are equivalent according to Definition 1. Let's next assume that the following code is used after both the original and the refactored code, respectively.

```
dim1.setSize(1,2);
dim2.setSize(2,3);
```

If the original and the refactored code were indeed equivalent, then we would expect *dim1* and *dim2* to have the same value at the end of both blocks of code. However, this is not the case. In the original code, *dim1* will have width 1 and height 2, whereas in the refactored code, *dim1* will have width 2 and height 3. This is due to the fact that, in the refactored code, *dim1* and *dim2* are aliases. Thus, when *dim2* has its size set to (2,3), this also affects *dim1*.

Intuitively, any aliases between live variables at the end of the original code should also be present at the end of the refactored code. For this purpose, we use the notation *aliasEquivClass(V)* which returns the set of all equivalence classes induced over the set of variables *V* by the aliasing relation. Notably, the aliasing equivalence relation must also hold over the static fields of the classes loaded by both programs.

EXAMPLE 4. For P_1 and P_2 , there are no aliases. However, for P_3 and P_4 we have:

```
aliasEquivClass(liveVar( $P_3$ )  $\cup$ 
staticFields(loadedClasses( $P_3$ ))) =  $\emptyset$ 
aliasEquivClass(liveVar( $P_4$ )  $\cup$ 
staticFields(loadedClasses( $P_4$ ))) = {{dim1, dim2}}
```

In P_4 , the aliasing relation induces one equivalence class, namely {*dim1*, *dim2*}.

Next, we complete Definition 1 by capturing the aliasing aspect.

DEFINITION 2 (ADDITION TO DEFINITION 1). In addition to Definition 1, two programs P_1 and P_2 are equivalent with respect to $\vec{i}$ iff:

$$(4) \forall v_1, v_2 \in \text{liveVar}(P_1) \cup \text{staticFields}(\text{loadedClasses}(P_1)). \\ \text{aliases}(P_1, v_1, v_2) \iff \text{aliases}(P_2, v_1, v_2)$$

where, given a program P , variables v_1 and v_2 are aliases, i.e., $\text{aliases}(P, v_1, v_2)$ holds if and only if they are in the same equivalence class induced by the aliasing relation, i.e., $\text{aliasEquivClass}(\{v_1, v_2\}) = \{v_1, v_2\}$.

We abuse the notation to use *staticFields* over a set of classes, rather than just one class. The objective is to return the union of all static fields defined in all the classes in the set of classes taken as argument.

3 Seeding of the component library

In this section, we describe how we populate the instruction library (also referred to as the component library) used by the symbolic approach starting from code hints. We'll refer to this library as

the CodeHints-library. This step is critical as, in order for the symbolic engine to succeed, the library must be as small as possible while containing all the instructions necessary for constructing the solution.

For illustration, let's go back to the running example in Figure 1 (with the corresponding Javadoc comment in Figure 2). As mentioned in Section 1, when building the corresponding component library for this example, we must find the necessary components that consume existing objects and create a *Calendar* object, so that we can call *method get* as suggested by the code hint. Besides adding components that allow us to generate the required objects, we also need components for setting their fields. In our example, we must call *calendar.setTime(date)* to set the calendar's date based on the existing date object.

Adding too many components to the library will make the code generation task infeasible. In particular, we should be able to differentiate between components that we can use (we have or we are able to generate all the necessary arguments and the current object for calling them), and those that we can't because we can't obtain some of the arguments and/or the current object. Adding the latter components to the library will significantly slow down the code generation process by adding infeasible programs to the search space. We address these challenges by dynamically building the component library for each refactoring such that it only contains components specific to that particular use case.

Throughout the seeding process, we keep track of the following sets: *consumable_objs* (inputs to the code to be refactored, which need to be consumed by the refactoring), *available_types* (types for which we either have consumable objects or the corresponding generators to create them) and *target_types* (types for which we must be able to generate objects). To start with, the *target_types* set contains the types of the original code's outputs.

For the running example, we start with: *consumable_objs* = {*date*}, *available_types* = {*Date*}, *target_types* = {*int*}. The objective of the seeding algorithm is to add components to the library that make use of the *available_types* to generate objects of *target_types*. At the same time, we want to consume the objects from the *consumable_objs* set, i.e., the inputs of the original code.

The seeding algorithm for the CodeHints-library is provided in Figure 5 and consists of three phases, which we discuss next.

Phase 1: Initialise with code hints. During the first phase, the initialisation, we add all the constants and instructions from the code hints to the library. For our running example, the hints in Figure 2 instruct us to add method "*int get(int field)*" from class *Calendar* and constant "*Calendar.HOUR_OF_DAY*" to our library. As a side note, finding the right constants is a well known challenge for program synthesis [2], and thus the subsequent synthesis process will always attempt to use the constants provided in the code hints before generating new ones.

Intuitively, we need to make sure that the Javadoc suggestions are realisable as captured by Definition 3, where *required_types(method)* refers to the types of the objects required to call *method* (i.e., the types corresponding to its arguments and current object). In our running example, *required_types(get)* = {*int*, *Calendar*} given that, in order to call *get*, we must provide an argument of type *int* and a current object of type *Calendar*. Method *get* is

```

/* @deprecated As of JDK version 1.1,
 * replaced by {@code contains(int, int)}.
 */
@Deprecated
public boolean inside(int X, int Y) { ... }

```

Figure 4: Javadoc hint example.

not realisable as the library doesn't contain any generator for `Calendar`. Consequently, `Calendar` is added to `target_types`, resulting in: `target_types = {int, Calendar}`.

DEFINITION 3 (REALISABLE METHOD). *Method i is realisable iff $\forall t \in \text{required_types}(i). t \in \text{available_types} \vee t$ is a primitive type.*

One challenge in this phase is interpreting the `@code` blocks inside `@deprecated` sections, which we attempt to parse as Java expressions. In particular, we parse each hint as if it were invoked in the context of the method to refactor, such that imports or the implicit `this` argument are considered during parsing.

In order to address the challenge that code hints are not always expressed as well-formed Java, we customised the GitHub Java parser to accept undeclared identifiers and type names as arguments. For instance, the Javadoc hint for deprecating `boolean inside(int X, int Y)` in Figure 4 suggests using `contains(int, int)`, which would normally cause a parsing error as the `int` type appears in the place of argument names. In our setting, we accept this hint as valid. There are still situations where our parser is too strict and fails to accept some of the code hints. Additionally, there are scenarios where, while the Javadoc does contain a useful code hint, it is not tagged accordingly with the `@code` tag.

Phase 2: Add generators for `target_types`. By generators we refer to constructors and any other methods returning objects of that particular type. In the algorithm in Figure 5, once we added a generator for a new type, we must add that type to `available_types`. Additionally, if the new generator consumes any objects from `consumable_objs`, we must remove them from the set.

For our running example, we must seed our component library with generators for `Calendar`. We first scan all the public constructors of class `Calendar` and all the public methods from class `Calendar` that return an object of type `Calendar`. We find the following four options:

- (1) **static** `Calendar getInstance()` – creates the object using the default time zone and locale.
- (2) **static** `Calendar getInstance(Locale)` – creates the object using the default time zone and specified locale.
- (3) **static** `Calendar getInstance(TimeZone)` – creates the object using the specified time zone and default locale.
- (4) **static** `Calendar getInstance(TimeZone, Locale)` – creates the object with the specified time zone and locale.

Out of the four methods, only the first one is realisable. Conversely, in order to call the second method, we would need to generate an object of type `Locale`, for which we don't have a corresponding component in the library, and the same applies to the last two methods. Thus, we only add the first method to the CodeHints-library.

Output: CodeHints-library

// **Phase 1: Initialise**

Add constants and instructions from the code hints to the library;

for each unrealisable instruction i in library do
 | `target_types = target_types  $\cup$  required_types( $i$ );`
end

// **Phase 2: Add generators for `target_types`**

for each $t \in \text{target_types}$ for which there is no generator do

 Add realisable generators for t to library;

`available_types = available_types  $\cup$  { $t$ };`

 Remove consumed objs from `consumable_objs`;

end

// **Phase 3: Add transformers for `target_types`**

while `consumable_objs` $\neq \emptyset$ do

 Add realisable transformers to the library for types in

`target_types` that consume objects from

`consumable_objs`;

 Remove consumed objs from `consumable_objs`;

end

Figure 5: Seeding algorithm for the CodeHints-library

Phase 3: Add transformers for `target_types`. The third and last phase adds transformers for the target types. By transformers we refer to methods that modify the value of an instance variable. Given our objective to generate objects for the `target_types` while consuming objects from `consumable_objs`, we prioritise transformers that consume such objects. For instance, for our running example there are 16 public transformers for the `Calendar` class. However, instead of adding all of them to the CodeHints-library, we prioritise those that consume the date object. There is only one such transformer `void setTime(Date date)`. Once we added any new transformers to the library, we remove the consumed objects from `consumable_objs`.

Notably, all the components that we add to the library must be accessible from the current location.

It may be the case that there is no realisable generator for a target type, or no transformers for the `target_types` that can consume all the `consumable_objs`. When we finish exploring all the available methods, we exit the corresponding loops in phases 2 and 3. As a consequence, the synthesiser may fail to generate a valid refactoring from the CodeHints-library. A possible future direction is allowing unrealisable generators and transformers to be added to the library and iterating phases 2 and 3 a given number of times (potentially until reaching a fixed point).

3.1 Types-library

For cases when code hints are not provided, in addition to the CodeHints-library, we also provide a component library that is populated based only on the type signature of the deprecated method. Next, we describe how this Types-library is being built.

For the CodeHints-library, the Javadoc code hints guidance results in the code hints being used to collect the *target_types* set during the initialisation phase. If we don't have access to code hints and we are seeding solely based on types, we start with the *target_types* set containing the types of the outputs of the original code. Phases 2 and 3 of the seeding algorithm are the same as those in Figure 5, where we attempt to add generators and transformers for the *target_types* to the library.

Compared to the CodeHints-library, the seeding process for the Types-library may end up missing critical types provided by the Javadoc code hints. For instance, in our running example, the `Calendar` class is only mentioned by the code hints and would not be included in the seeding of the Types-library. One possibility for adding `Calendar` to the target types without considering the Javadoc hints, would be to add all the classes from the `java.util` package, which contains `Date`. However, doing so would result in a very large component library, most likely outside the capabilities of existing synthesis techniques.

4 Design and implementation choices

4.1 Abstract classes and interfaces

The programs that need to be checked for equivalence may refer to abstract classes and interfaces. These are by default instantiated using the mocking framework Mockito. Alternatively, we also curate a list of explicit constructors of subclasses to be used in favour of Mockito mocks for certain types, and users have the option to extend this list with a custom configuration. This is particularly useful when attempting to avoid constructors that are not amenable to fuzzing, such as collection constructors that take an integer capacity argument, where an unlucky fuzzed input might lead to an out of memory error.

4.2 Observable state

Our equivalence check uses the Java Reflection API to observe the side effects of programs, such as assigning new values to fields. As a consequence, we cannot observe effects that are not visible through this API, such as I/O operations or variables maintained in native code only. I/O operations could be recorded using additional bytecode instrumentation in future work, but for the scope of our current implementation, if programs only differ in such side effects, our engine is prone to producing a refactoring that it's not fully equivalent to the original.

4.3 Instrumentation and isolation

We use reflection to invoke the original method to refactor with fuzzed inputs, as well as our synthesised refactoring candidates. This allows us to dynamically invoke new candidates without the need for compilation. Regarding static fields, Java allows to load the same class in different class loaders, which creates independent copies of its static fields.

In order to fully isolate the program state during the execution of the original and refactored code from each other, we load all involved classes in separate class loaders. These class loaders are disposed immediately after the current set of fuzzed inputs are executed, and new class loaders are created for the next inputs.

Classes that do not maintain a static state are loaded in a shared parent class loader to improve performance.

The OpenJDK Java virtual machine implementation does not allow us to load classes contained in the `java.lang` package in such an isolated class loader. For such classes we cannot apply refactorings that depend on static fields, as they cannot be reset and will retain the state of previous executions. Instead of observing the effect of one isolated refactoring candidate, we would observe their accumulated effect, which would be incorrect. For the scope of our experiments this did not pose a problem, since most of the affected classes in the `java.lang` package do not maintain a static state, and the ones who do were irrelevant to our refactorings.

4.4 Checking aliasing

While in Section 2.3, we discussed aliasing preservation in terms of preserving the equivalence classes induced by the aliasing relation, in our implementation we enforce a simpler but stricter than necessary check. In particular, all the objects that are being referenced during the code's execution are assigned increasing symbolic identifiers. Then, we enforce aliasing preservation (condition (4) in Definition 2), by checking that, for all variables defined by both the original and the refactored code, the objects they reference have the same symbolic id.

For illustration, in the original code in Example 3, the object of type `Dimension` referenced by variable *dim1* is assigned symbolic value s_1 , whereas the object referenced by *dim2* is assigned symbolic value s_2 . In a similar manner, the object referenced by both *dim1* and *dim2* in the refactored code is assigned value s_1 . Then, the aliasing check fails as, in the original code, the object referenced by *dim2* has symbolic value s_2 , whereas, in the refactored code, the object referenced by *dim2* has symbolic value s_1 . While this is a stronger than needed requirement, it was sufficient for our experiments. In particular, we didn't find any benchmark where a sound refactoring was rejected due to it.

5 Experimental evaluation

We implemented the refactoring generation techniques in two tools called `REFSYM` and `REFNEURAL`, corresponding to the symbolic and the neural approach, respectively (available together with all the experiments at <https://github.com/pkesseli/refactoring-synthesis/tree/hanliang/dev>).

5.1 Experimental setup

Our benchmark suite contains 236 out of 392 deprecated methods in the Oracle JDK 15 Deprecated API documentation [41]. We included all the deprecated methods except those that were deprecated without a replacement (e.g. `java.rmi.registry.RegistryHandler.registryImpl(int)`), methods inaccessible by non-JCL classes (i.e., all the `finalize` methods, which are called by the garbage collector, not user code, meaning that we can't modify their calls), methods that only perform I/O (as our equivalence check doesn't include I/O side effects), and native methods.

For `REFNEURAL`, we use Claude 2.1 and Claude 3, denoted as `REFNEURAL-Claude2.1` and `REFNEURAL-Claude3`, respectively. Claude 2.1 and Claude 3 are proprietary LLMs likely to be very large (1T+

parameters), which have been shown to be among the highest performing on coding tasks [10, 21, 27, 36]. We accessed the LLMs via Amazon Bedrock. To make our results more deterministic, we use a lower temperature (i.e. less random) of 0.2. We also repeated the experiments three times, and the results are summarized as averages. For all engines, we bound the search by at most 500 inputs and 5 minutes per verification phase, and at most 2 minutes per synthesis phase.

Experiments were performed on an Ubuntu 22.04 x64 operating system running in a laptop with 16 GB RAM and 11th Gen Intel Core i7-11850H at 2.50 GHz. The JVM used was Oracle JDK 15.02.

For REF_{SYM}, if code hints are present, we start with the CodeHints-library and fall back to the Types-library if the synthesis phase (as described in Section 2.1) times out for the CodeHints-library; if no code hints exist then we use the Types-library.

5.2 Results

Configuration	✓	✗	⚡	%	∅ runtime (s)
REF _{SYM}	7	82	16	6	213.7
REF _{NEURAL} -Claude2.1	12	91	2	11	197.14
REF _{NEURAL} -Claude3	10	93	2	9	203.32
Best Virtual Engine	15	77	13	14	195.9
REF _{SYM} (CH)	104	19	8	79	220.8
REF _{NEURAL} -Claude2.1 (CH)	93.3	36.7	1	71	212.19
REF _{NEURAL} -Claude3 (CH)	94	36	1	72	217.9
Best Virtual Engine (CH)	107	13	11	82	210.28

Table 1: Experimental results for all benchmarks.

Table 1 provides an overview of the number of sound refactorings (✓), missed refactorings (✗), unsound refactorings (⚡), the percentage of sound refactorings (%) produced per configuration, as well as the average runtime per refactoring. For the symbolic engine, refactorings are guaranteed to compile so missed refactorings are those that failed our equivalence check, whereas for the neural engine, they either didn't compile or failed the equivalence check. Unsound refactorings are those that passed the equivalence check, but were manually detected by us as not being equivalent to the original code. We describe the reasons why this can happen later in this section. The average runtime includes both the time to generate a refactoring and to verify it.

We split our dataset into benchmarks where code hints could be extracted from the Javadoc, and benchmarks without code hints. The first four rows provide the results for benchmarks without code hints, whereas the last four (marked with "(CH)") show the results for those benchmarks where code hints were present. The rows denoted by "Best Virtual Engine" count all benchmarks (with and without code hints, respectively) solved by at least one engine.

The results support our hypothesis that code hints are very valuable for automating the refactoring process. For all engines, benchmarks with code hints have a considerably higher success rate than those without code hints, with the best virtual engine solving 82% of the benchmarks with code hints. In the absence of code hints, all the engines are struggling, solving only a few benchmarks.

To understand the discrepancy between the number of missed/unsound refactorings with the without code hints, we next investigate the main reasons for such refactorings.

Missed refactorings. For the symbolic engine, the majority of the missed refactorings were due to fuzzing timeouts, which are likely caused by the component libraries missing some instructions needed in the synthesis phase. As explained in Section 3, one option for increasing the size of our libraries is allowing unrealisable generators and transformers to be added.

For the neural approach, we could only observe that the LLMs were unable to generate the refactoring from the provided prompt.

For both engines, when refactoring methods that eventually run native code, we encountered crashes of the verifier, which, in order to be conservative, we counted as overall missed refactorings. The reason for this is the fact that we use reflection to produce counterexamples, and some may violate internal invariants. If they execute methods that eventually call native code, this can lead to the entire JVM crashing rather than e.g. throwing an exception.

Unsound refactorings. In several cases, refactorings that are not equivalent to the original code managed to pass our verifier. While we expected to run into this problem because of the nature of fuzzing (the fuzzer may miss counterexamples that distinguish the behaviour of the original from the refactored code), we also encountered other problems:

Unobservable state (discussed in Section 4.2): I/O operations, static state in the boot class loader and native methods are not observable by our equivalence predicate implementation, since we rely on reflection to examine objects inside the Java runtime. As a consequence, we cannot distinguish programs that only differ in these aspects.

Abstract methods: There are classes in the JCL without any existing implementation (e.g. `javax.swing.InputVerifier`), and thus no concrete method against which to verify the equivalence between the original and the refactored code. For those, the symbolic engine will generate a no-op. We've been very conservative here, and counted this scenario as "unsound" because it doesn't match human intent for the deprecated method.

Insufficient counterexamples: Our fuzzing-based equivalence check is inherently incomplete, and for some benchmarks we do not explore sufficient counterexamples to identify unsound candidates. An example is `javax.swing.JViewport#isBackingStoreEnabled`, where only a single input out of the $2^{32} - 1$ possible input values will trigger an alternate code path. As a second exemplar, method `java.rmi.server.RMIClassLoader#loadClass` accepts a string as an input, but will throw when given any string that is not a valid class name. It is very unlikely that our fuzzer will randomly produce a string that matches a valid class name.

Discussion. The symbolic engine benefits significantly from code hints, as hints seed the component library, making it less likely that needed instructions are missing. For the neural engine, we hypothesise that, by providing additional context to the LLM, code hints are aiding the code generation task.

Benchmarks with code hints also have fewer unsound results. Our hypothesis is that both engines are much more likely to generate the expected refactoring before generating other candidate refactorings that may trick our equivalence checker.

5.3 Research questions

(RQ1) Can the refactoring of deprecated Java APIs be automated? Yes, the refactoring of deprecated Java APIs can be automated if code hints are added to the JDK, evidenced by the 82% success rate for those benchmarks. This automation could greatly assist engineers with the code migration task.

(RQ2) Do code hints help the generation of refactorings? The performance of REF_{SYM}, REF_{NEURAL}-Claude2.1 and REF_{NEURAL}-Claude3 improves considerably when code hints are present. All of them barely manage to solve any benchmark without code hints. When code hints are present, all engines solve at least 71% of benchmarks, with the best virtual engine solving 82%. For the symbolic engine, they enable the effective seeding of the instruction library, whereas for the neural engine, they provide additional context to the LLM. In summary, without code hints, Java's complexity makes the refactoring of deprecated Java APIs very hard for both for the symbolic and neural engines.

(RQ3) How do the symbolic and the neural approach compare against each other? The symbolic and the neural engines have very similar performance (where the symbolic engine has a smaller computational cost).

When code hints are present, as shown in Table 1, the symbolic approach does slightly better than both REF_{NEURAL}-Claude2.1 and REF_{NEURAL}-Claude3. This supports the intuition that, if there is enough information about the solution to effectively prune the solution space (in our case, when code hints are present), then the symbolic approach works well. One example of a benchmark that was solved by REF_{SYM} but where both REF_{NEURAL}-Claude2.1 and REF_{NEURAL}-Claude3 failed to find a solution is the running example in Figure 1. In all our runs, the LLMs failed to generate `calendar.setTime(date)`.

When code hints are not present, the neural engine does marginally better than the symbolic approach. Compared to the symbolic engine, the neural one is able to provide correct refactorings for deprecated concurrency primitives (e.g., `weakCompareAndSet` in `java.util.concurrent.atomic.AtomicReference`). These methods call native methods, for which we cannot observe side-effects (Section 4.2). Consequently, these are not captured by the counterexamples returned by the fuzzer, and are thus not taken into consideration by the symbolic engine during the synthesis phase. For the neural approach, the counterexamples are less critical, and the LLM can obviously generate the correct code even though they are incomplete.

The neural engine is also able to handle benchmarks that require special constants, such as `java.net.URLDecoder.decode(s)` described in the introduction.

6 Threats to validity

Selection of benchmarks. All our benchmarks are Java methods deprecated in the Oracle JDK 15. The JDK is extremely well known

to Java developers, and a lot of Java application code evolves similarly. However, our claims may not extend to other programming languages.

Quality of refactorings. Refactorings need to result in code that remains understandable and maintainable. It is difficult to assess objectively how well our technique does with respect to this subjective goal. This threatens our claim that refactoring of deprecated Java APIs be automated.

We manually inspected the refactorings obtained with both engines and found them to represent sensible transformations.

Efficiency and scalability of the program synthesiser. We apply program synthesis and fuzzing. This implies that our broader claim is threatened by scalability limits of these techniques. While for the majority of our experiments the synthesiser was able to find a solution, there were a few cases where it timed out, either because it could not generate a candidate, or it could not verify it. Component libraries with a diverse range of component sizes may help mitigate this effect.

Prompt engineering for Claude. In our current experiments, we engineered prompts for the Claude LLMs. While we made best efforts to follow the official prompt engineering guidelines², which presents prompt design techniques to improve model performance, there might be better ways of composing it. This might invalidate our claim that the symbolic and neural techniques deliver roughly the same performance.

7 Related Works

Program refactoring. We first discuss works on the refactoring of deprecated instances. The work of Perkins directly replaces calls to deprecated methods by their bodies [45], but we argue this conflicts with the intent of language designers. Moreover, it can introduce concurrency bugs as inlining calls to deprecated methods can cause undesirable effects if the original function was synchronised. While it is not possible for two invocations on the same object of the synchronised original method to interleave, this is not guaranteed after inlining the method's body. Notably, the authors of [47] show how refactorings in concurrent programs can inadvertently introduce concurrency issues by enabling new interactions between parallel threads.

A related class of techniques aim to adapt APIs after library updates. Such techniques automatically identify change rules linking different library releases [22, 55]. Conversely to our work, a change rule describes a match between methods existing in the old release, but which have been removed or deprecated in the new one, and replacement methods in the new release. However, they do not provide the actual refactored code. In [33], Lee et al. address the problem of outdated APIs in documentation references. Their insight is that API updates in documentation can be derived from API implementation changes between code revisions. Conversely, we are looking at code changes, rather than documentation. Other works focus on automatically updating API usages for Android apps based on examples of how other developers evolved their apps for the same changes [11, 19]. Furthermore, [20] improves

²<https://docs.anthropic.com/en/docs/prompt-engineering>

on [19] by using a data-flow analysis to resolve the values used as API arguments and variable name denormalization to improve the readability of the updated code. As opposed to these works, which rely on examples of similar fixes, our approach uses Javadoc code hints.

Several rule-based source-to-source transformation systems exist that allow users to define transformation rules based on the program's syntax [5, 54]. Unlike our approach, these systems require users to manually specify the transformation rules. In the context of code translation, a few works have explored the automatic generation of such rules using static analysis, but these efforts have been specialised to C-to-Rust translation [9, 61].

Search-based approaches to automating the task of software refactoring, based on the concept of treating object-oriented design as a combinatorial optimisation problem, have also been proposed [38, 39]. They usually make use of techniques such as simulated annealing, genetic algorithms and multiple ascent hill-climbing. While our technique is search based, the search is guided by types, code hints, as well as counterexamples, discovered by testing.

Closer to our work, several refactoring techniques make explicit use of some form of semantic information. Khatchadourian et al. use a type inference algorithm to automatically transform legacy Java code (pre Java 1.5) to use the enum construct [31]. Other automated refactoring techniques aim to transform programs to use a particular design pattern [4, 25]. Steimann et al. present Constraint-Based Refactoring [49–51], where given well-formedness logical rules about the program are translated into constraints that are then solved to assist the refactoring. Fuhrer et al. implement a type constraint system to introduce missing type parameters in uses of generic classes [15] and to introduce generic type parameters into classes that do not provide a generic interfaces despite being used in multiple type contexts [32]. Kataoka et al. use program invariants (found by the dynamic tool Daikon) to infer whether specific refactorings are applicable [29]. Finding invariants is notoriously difficult. Moreover, the technique is limited to a small number of refactorings and does not include the elimination of deprecated instances. In [18], Gyori et al. present the tool LAMBDAFICATOR, which automates two pattern-based refactorings. The first refactoring converts anonymous inner classes to lambda expressions. The second refactoring converts for loops that iterate over Collections to functional operations that use lambda expressions. The LAMBDAFICATOR tool [14] is available as a NetBeans branch.

Our techniques does not expect any precomputed information such as logical well-formedness properties or invariants. Instead, we make use of information that is already available in the original program and Javadoc, namely code hints and type information. Moreover, we explore the space of all potential candidate programs by combining techniques from type-directed, component-based and counterexample-guided inductive synthesis.

Symbolic program synthesis. CEGIS-based approaches to program synthesis have been previously used for program transformations such as superoptimisation and deobfuscation [26]. In [6], David et al. present an automatic refactoring tool that transforms Java with external iteration over collections into code that uses Streams. Their approach makes use of formal verification to check the correctness of a refactoring. Cheung et al. describe a system

that transforms fragments of application logic into SQL queries [3] by using a CEGIS-based synthesiser to generate invariants and postconditions validating their transformations (a similar approach is presented in [23]). While our approach is also CEGIS-based, we are guided by code hints and types to efficiently prune the search space.

Type information has been extensively used in program synthesis to guide the search for a solution [13, 30, 35, 42, 57]. In our work, we combine type information with code hints. Another direction that inspired us is that of component-based synthesis [12, 16, 26], where the target program is generated by composing components from a library. Similarly to these approaches, we use a library of components for our program generation approach. However, our technique uses information about types and code hints to build the component library, which is specific to each refactoring.

Large Language Models. LLMs have been used successfully for code generation tasks, with applications ranging from code completions [7, 8, 24, 37], translations [44, 52] to repository-level generation [59] and general software engineering tasks [58]. Those models are typically pre-trained on vast amounts of data, fine-tuned for specific tasks, and require advanced prompt engineering [28]. Among the well-known LLMs, GPT-4 [40], Claude2.1 [1], Claude3 [1], LLaMa [53] are general-purpose models covering a diverse set of language-related applications; there are also models pre-trained/fine-tuned specifically for code generation and programming tasks, such as CodeLLaMa2 [46], StarCoder2 [34], DeepSeek-Coder [17], and GrammarT5 [63]. In this work, we have applied the Claude family of LLMs to our refactoring task, with the goal of investigating how much it benefits from Javadoc code hints.

8 Conclusions

In this paper, we investigated the benefits of Javadoc code hints when refactoring deprecated methods. For this purpose, we designed and implemented a symbolic and a neural automatic program refactoring techniques that eliminate uses of deprecated methods. In our experiments, both engines demonstrate strong performance when code hints were present, and fare much worse otherwise, leading us to conclude that code hints can boost the automation of refactoring code that uses deprecated Java APIs.

Acknowledgements. Cristina David was supported by the Royal Society University Research Fellowship URF\R\221031.

References

- [1] [n. d.]. Claude. <https://www.anthropic.com/index/introducing-claude>.
- [2] Alessandro Abate, Cristina David, Pascal Kesseli, Daniel Kroening, and Elizabeth Polgreen. 2018. Counterexample Guided Inductive Synthesis Modulo Theories. In *Computer Aided Verification - 30th International Conference, CAV 2018, Held as Part of the Federated Logic Conference, FloC 2018, Oxford, UK, July 14–17, 2018, Proceedings, Part I (Lecture Notes in Computer Science, Vol. 10981)*, Hana Chockler and Georg Weissenbacher (Eds.). Springer, 270–288. doi:10.1007/978-3-319-96145-3_15
- [3] Alvin Cheung, Armando Solar-Lezama, and Samuel Madden. 2013. Optimizing database-backed applications with query synthesis. In *Conference on Programming Language Design and Implementation, PLDI*. 3–14.
- [4] Aikaterini Christopoulou, E.A. Giakoumakis, Vassilis E. Zafeiris, and Soukara Vasiliki. 2012. Automated refactoring to the Strategy design pattern. *Information and Software Technology* 54, 11 (2012), 1202 – 1214.
- [5] James R Cordy, Thomas R Dean, Andrew J Malton, and Kevin A Schneider. 2002. Source transformation in software engineering using the TXL transformation system. *Information and Software Technology* 44, 13 (2002), 827 – 837.

- [6] Cristina David, Pascal Kesseli, and Daniel Kroening. 2017. Kayak: Safe Semantic Refactoring to Java Streams. *CoRR* abs/1712.07388 (2017). arXiv:1712.07388 <https://arxiv.org/abs/1712.07388>
- [7] Yangruibo Ding, Zijian Wang, Wasi Ahmad, Murali Krishna Ramanathan, Ramesh Nallapati, Parminder Bhatia, Dan Roth, and Bing Xiang. 2024. CoCoMIC: Code completion by jointly modeling in-file and cross-file context. In *LREC-COLING 2024*. <https://www.amazon.science/publications/cocomic-code-completion-by-jointly-modeling-in-file-and-cross-file-context>
- [8] Xueying Du, Mingwei Liu, Kaixin Wang, Hanlin Wang, Junwei Liu, Yixuan Chen, Jiayi Feng, Chaofeng Sha, Xin Peng, and Yiling Lou. 2024. Evaluating Large Language Models in Class-Level Code Generation. In *Proceedings of the IEEE/ACM 46th International Conference on Software Engineering (ICSE '24)*. Association for Computing Machinery, New York, NY, USA, Article 81, 13 pages. doi:10.1145/3597503.3639219
- [9] Mehmet Emre, Ryan Schroeder, Kyle Dewey, and Ben Hardekopf. 2021. Translating C to safer Rust. *Proc. ACM Program. Lang.* 5, OOPSLA, Article 121 (Oct. 2021), 29 pages. doi:10.1145/3485498
- [10] Hasan Ferit Eniser, Hanliang Zhang, Cristina David, Meng Wang, Maria Christakis, Brandon Paulsen, Joey Dodds, and Daniel Kroening. 2024. Towards Translating Real-World Code with LLMs: A Study of Translating to Rust. *CoRR* abs/2405.11514 (2024). doi:10.48550/ARXIV.2405.11514 arXiv:2405.11514
- [11] Mattia Fazzini, Qi Xin, and Alessandro Orso. 2019. Automated API-usage update for Android apps. In *Proceedings of the 28th ACM SIGSOFT International Symposium on Software Testing and Analysis, ISSTA 2019, Beijing, China, July 15-19, 2019*, Dongmei Zhang and Anders Møller (Eds.). ACM, 204–215. doi:10.1145/3293882.3330571
- [12] Yu Feng, Ruben Martins, Yuepeng Wang, Isil Dillig, and Thomas W. Reps. 2017. Component-based synthesis for complex APIs. In *Proceedings of the 44th ACM SIGPLAN Symposium on Principles of Programming Languages, POPL 2017, Paris, France, January 18-20, 2017*, Giuseppe Castagna and Andrew D. Gordon (Eds.). ACM, 599–612. doi:10.1145/3009837.3009851
- [13] John K. Feser, Swarat Chaudhuri, and Isil Dillig. 2015. Synthesizing data structure transformations from input-output examples. In *Proceedings of the 36th ACM SIGPLAN Conference on Programming Language Design and Implementation, Portland, OR, USA, June 15-17, 2015*, David Grove and Stephen M. Blackburn (Eds.). ACM, 229–239. doi:10.1145/2737924.2737977
- [14] Lyle Franklin, Alex Gyori, Jan Lahoda, and Danny Dig. 2013. LAMBDAFICATOR: from imperative to functional programming through automated refactoring. In *35th International Conference on Software Engineering, ICSE '13, San Francisco, CA, USA, May 18-26, 2013*. 1287–1290. doi:10.1109/ICSE.2013.6606699
- [15] Robert M. Fuhrer, Frank Tip, Adam Kiezun, Julian Dolby, and Markus Keller. 2005. Efficiently Refactoring Java Applications to Use Generic Libraries. In *ECOOP 2005 - Object-Oriented Programming, 19th European Conference, Glasgow, UK, July 25-29, 2005*, *Proceedings*. 71–96. doi:10.1007/11531142_4
- [16] Sumit Gulwani, Susmit Jha, Ashish Tiwari, and Ramarathnam Venkatesan. 2011. Synthesis of loop-free programs. In *Proceedings of the 32nd ACM SIGPLAN Conference on Programming Language Design and Implementation, PLDI 2011, San Jose, CA, USA, June 4-8, 2011*, Mary W. Hall and David A. Padua (Eds.). ACM, 62–73. doi:10.1145/1993498.1993506
- [17] Daya Guo, Qihao Zhu, Dejian Yang, Zhenda Xie, Kai Dong, Wentao Zhang, Guanting Chen, Xiao Bi, Y. Wu, Y. K. Li, Fuli Luo, Yingfei Xiong, and Wenzheng Liang. 2024. DeepSeek-Coder: When the Large Language Model Meets Programming – The Rise of Code Intelligence. arXiv:2401.14196 [cs.SE]
- [18] Alex Gyori, Lyle Franklin, Danny Dig, and Jan Lahoda. 2013. Crossing the gap from imperative to functional programming through refactoring.. In *ESEC/SIGSOFT FSE*. ACM, 543–553.
- [19] Stefanus A. Haryono, Ferdian Thung, Hong Jin Kang, Lucas Serrano, Gilles Muller, Julia Lawall, David Lo, and Lingxiao Jiang. 2020. Automatic Android Deprecated-API Usage Update by Learning from Single Updated Example. In *ICPC'20: 28th International Conference on Program Comprehension, Seoul, Republic of Korea, July 13-15, 2020*. ACM, 401–405. doi:10.1145/3387904.3389285
- [20] Stefanus A. Haryono, Ferdian Thung, David Lo, Lingxiao Jiang, Julia Lawall, Hong Jin Kang, Lucas Serrano, and Gilles Muller. 2022. AndroEvolve: automated Android API update with data flow analysis and variable denormalization. *Empir. Softw. Eng.* 27, 3 (2022), 73. doi:10.1007/S10664-021-10096-0
- [21] Wenpin Hou and Zhicheng Ji. 2024. Comparing large language models and human programmers for generating programming code. arXiv:2403.00894 [cs.SE] <https://arxiv.org/abs/2403.00894>
- [22] Kaifeng Huang, Bihuan Chen, Linghao Pan, Shuai Wu, and Xin Peng. 2021. REPFINDER: Finding Replacements for Missing APIs in Library Update. In *36th IEEE/ACM International Conference on Automated Software Engineering, ASE 2021, Melbourne, Australia, November 15-19, 2021*. IEEE, 266–278. doi:10.1109/ASE51524.2021.9678905
- [23] Ming-Yee Iu, Emmanuel Cecchet, and Willy Zwaenepoel. 2010. JReq: Database Queries in Imperative Languages. In *Compiler Construction (CC)*. 84–103.
- [24] M. Izadi, J. Katzy, T. van Dam, M. Otten, R. Popescu, and A. van Deursen. 2024. Language Models for Code Completion: A Practical Evaluation. In *2024 IEEE/ACM 46th International Conference on Software Engineering (ICSE)*. IEEE Computer Society, Los Alamitos, CA, USA, 956–968. <https://doi.ieeecomputersociety.org/>
- [25] Sang-Uk Jeon, Joon-Sang Lee, and Doo-Hwan Bae. 2002. An automated refactoring approach to design pattern-based program transformations in Java programs. In *Asia-Pacific Software Engineering Conference (APSEC)*. 337–345.
- [26] Susmit Jha, Sumit Gulwani, Sanjit A. Seshia, and Ashish Tiwari. 2010. Oracle-guided component-based program synthesis. In *Proceedings of the 32nd ACM/IEEE International Conference on Software Engineering - Volume 1, ICSE 2010, Cape Town, South Africa, 1-8 May 2010*, Jeff Kramer, Judith Bishop, Premkumar T. Devanbu, and Sebastián Uchitel (Eds.). ACM, 215–224. doi:10.1145/1806799.1806833
- [27] Juyong Jiang, Fan Wang, Jiasi Shen, Sungju Kim, and Sunghun Kim. 2024. A Survey on Large Language Models for Code Generation. arXiv:2406.00515 [cs.CL] <https://arxiv.org/abs/2406.00515>
- [28] Juyong Jiang, Fan Wang, Jiasi Shen, Sungju Kim, and Sunghun Kim. 2024. A Survey on Large Language Models for Code Generation. arXiv:2406.00515 [cs.CL]
- [29] Yoshio Kataoka, David Notkin, Michael D. Ernst, and William G. Griswold. 2001. Automated Support for Program Refactoring Using Invariants. In *Proceedings of the IEEE International Conference on Software Maintenance (ICSM '01)*. IEEE Computer Society.
- [30] Susumu Katayama. 2005. Systematic search for lambda expressions. In *Revised Selected Papers from the Sixth Symposium on Trends in Functional Programming, TFP 2005, Tallinn, Estonia, 23-24 September 2005 (Trends in Functional Programming, Vol. 6)*, Marko C. J. D. van Eekelen (Ed.). Intellect, 111–126.
- [31] Raffi Khachatourian, Jason Sawin, and Atanas Rountev. 2007. Automated Refactoring of Legacy Java Software to Enumerated Types. In *Software Maintenance, 2007. ICSM 2007. IEEE International Conference on*. 224–233.
- [32] Adam Kiezun, Michael D. Ernst, Frank Tip, and Robert M. Fuhrer. 2007. Refactoring for Parameterizing Java Classes. In *29th International Conference on Software Engineering (ICSE 2007)*, Minneapolis, MN, USA, May 20-26, 2007. 437–446. doi:10.1109/ICSE.2007.70
- [33] Seonah Lee, Rongxin Wu, Shing-Chi Cheung, and Sungwon Kang. 2021. Automatic Detection and Update Suggestion for Outdated API Names in Documentation. *IEEE Trans. Software Eng.* 47, 4 (2021), 653–675. doi:10.1109/TSE.2019.2901459
- [34] Anton Lozhkov, Raymond Li, Loubna Ben Allal, Federico Scholok, Joel Lamy-Poirier, Nouamane Tazi, Ao Tang, Dmytro Pykhtar, Jiawei Liu, Yuxiang Wei, Tianyang Liu, Max Tian, Denis Kocetkov, Arthur Zucker, Younes Belkada, Zijian Wang, Qian Liu, Dmitry Abulkhanov, Indraneil Paul, Zhuang Li, Wen-Ding Li, Megan Risdal, Jia Li, Jian Zhu, Terry Yue Zhuo, Evgenii Zheltonozhskii, Nii Osae Osae Dade, Wenhao Yu, Lucas Krauß, Naman Jain, Yixuan Su, Xuanli He, Manan Dey, Edoardo Abati, Yekun Chai, Niklas Muennighoff, Xiangru Tang, Muh-tasham Oblokulov, Christopher Akiki, Marc Marone, Chenghao Mou, Mayank Mishra, Alex Gu, Binyuan Hui, Tri Dao, Armel Zebaze, Olivier Dehaene, Nicolas Patry, Canwen Xu, Julian McAuley, Han Hu, Torsten Scholak, Sebastian Paquet, Jennifer Robinson, Carolyn Jane Anderson, Nicolas Chapados, Mostofa Patwary, Nima Tajbakhsh, Yacine Jernite, Carlos Muñoz Ferrandis, Lingming Zhang, Sean Hughes, Thomas Wolf, Arjun Guha, Leandro von Werra, and Harm de Vries. 2024. StarCoder 2 and The Stack v2: The Next Generation. arXiv:2402.19173 [cs.SE]
- [35] Justin Lubin, Nick Collins, Cyrus Omar, and Ravi Chugh. 2020. Program sketching with live bidirectional evaluation. *Proc. ACM Program. Lang.* 4, ICFP (2020), 109:1–109:29. doi:10.1145/3408991
- [36] Lincoln Murr, Morgan Grainger, and David Gao. 2023. Testing LLMs on Code Generation with Varying Levels of Prompt Specificity. arXiv:2311.07599 [cs.SE] <https://arxiv.org/abs/2311.07599>
- [37] Ansong Ni, Srinu Iyer, Dragomir Radev, Ves Stoyanov, Wen-tau Yih, Sida I. Wang, and Xi Victoria Lin. 2023. LEVER: Learning to verify language-to-code generation with execution. In *Proceedings of the 40th International Conference on Machine Learning (Honolulu, Hawaii, USA), (ICML '23)*. JMLR.org, Article 1086, 23 pages.
- [38] M. O’Keeffe and M.O. Cinnéide. 2008. Search-based refactoring: an empirical study. *Journal of Software Maintenance and Evolution: Research and Practice* 20, 5 (2008), 345–364.
- [39] M. O’Keeffe and M.O. Cinnéide. 2008. Search-based refactoring for software maintenance. *Journal of Systems and Software* 81, 4 (2008), 502 – 516.
- [40] OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, Red Avila, Igor Babuschkin, Suchir Balaji, Valerie Balcom, Paul Baltescu, Haiming Bao, Mohammad Bavarian, Jeff Belgum, Irwan Bello, Jake Berdine, Gabriel Bernadett-Shapiro, Christopher Berner, Lenny Bogdonoff, Oleg Boiko, Madeline Boyd, Anna-Luisa Brakman, Greg Brockman, Tim Brooks, Miles Brundage, Kevin Button, Trevor Cai, Rosie Campbell, Andrew Cann, Britany Carey, Chelsea Carlson, Rory Carmichael, Brooke Chan, Che Chang, Fotis Chantzis, Derek Chen, Sully Chen, Ruby Chen, Jason Chen, Mark Chen, Ben Chess, Chester Cho, Casey Chu, Hyung Won Chung, Dave Cummings, Jeremiah Currier, Yunxing Dai, Cory Decareaux, Thomas Degry, Noah Deutsch, Damien Deville, Arka Dhar, David Dohan, Steve Dowling, Sheila Dunning, Adrien Ecoffet, Atty Eleti, Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix, Simón Posada Fishman, Juston Forte, Isabella Fulford, Leo Gao, Elie Georges, Christian Gibson, Vik Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-Lopes, Jonathan Gordon, Morgan Grafstein, Scott Gray, Ryan Greene, Joshua Gross, Shixiang Shane Gu, Yufei Guo, Chris Hallacy, Jesse Han, Jeff Harris, Yuchen He, Mike Heaton,

- Johannes Heidecke, Chris Hesse, Alan Hickey, Wade Hickey, Peter Hoeschele, Brandon Houghton, Kenny Hsu, Shengli Hu, Xin Hu, Joost Huizinga, Shantanu Jain, Shawn Jain, Joanne Jang, Angela Jiang, Roger Jiang, Haozhun Jin, Denny Jin, Shino Jomoto, Billie Jonn, Heewoo Jun, Tomer Kaftan, Lukasz Kaiser, Ali Kamali, Ingmar Kanitscheider, Nitish Shirish Keskar, Tabarak Khan, Logan Kilpatrick, Jong Wook Kim, Christina Kim, Yongjik Kim, Jan Hendrik Kirchner, Jamie Kiro, Matt Knight, Daniel Kokotajlo, Lukasz Kondraciuk, Andrew Kondrich, Aris Konstantinidis, Kyle Kosic, Gretchen Krueger, Vishal Kuo, Michael Lampe, Ikai Lan, Teddy Lee, Jan Leike, Jade Leung, Daniel Levy, Chak Ming Li, Rachel Lim, Molly Lin, Stephanie Lin, Mateusz Litwin, Theresa Lopez, Ryan Lowe, Patricia Lue, Anna Makanju, Kim Malfacini, Sam Manning, Todor Markov, Yaniv Markovski, Bianca Martin, Katie Mayer, Andrew Mayne, Bob McGrew, Scott Mayer McKinney, Christine McLeavey, Paul McMillan, Jake McNeil, David Medina, Aalok Mehta, Jacob Menick, Luke Metz, Andrey Mishchenko, Pamela Mishkin, Vinnie Monaco, Evan Morikawa, Daniel Mossing, Tong Mu, Mira Murati, Oleg Murk, David Mély, Ashvin Nair, Reiichiro Nakano, Rajeev Nayak, Arvind Neelakantan, Richard Ngo, Hyeonwoo Noh, Long Ouyang, Cullen O'Keefe, Jakub Pachocki, Alex Paino, Joe Palermo, Ashley Pantuliano, Giamattista Parascandolo, Joel Parish, Emy Parparita, Alex Passos, Mikhail Pavlov, Andrew Peng, Adam Perelman, Filipe de Avila Belbute Peres, Michael Petrov, Henrique Ponde de Oliveira Pinto, Michael, Pokorny, Michelle Pokrass, Vitchyr H. Pong, Tolly Powell, Alethea Power, Boris Power, Elizabeth Proehl, Raul Puri, Alec Radford, Jack Rae, Aditya Ramesh, Cameron Raymond, Francis Real, Kendra Rimbach, Carl Ross, Bob Rotsted, Henri Roussez, Nick Ryder, Mario Saltarelli, Ted Sanders, Shibani Santurkar, Girish Sastry, Heather Schmidt, David Schnurr, John Schulman, Daniel Selsam, Kyla Sheppard, Toki Sherbakov, Jessica Shieh, Sarah Shoker, Pranav Shyam, Szymon Sidor, Eric Sigler, Maddie Simens, Jordan Sitkin, Katarina Slama, Ian Sohl, Benjamin Sokolowsky, Yang Song, Natalie Staudacher, Felipe Petroski Such, Natalie Summers, Ilya Sutskever, Jie Tang, Nikolas Tezak, Madeleine B. Thompson, Phil Tillet, Amin Tootoonchian, Elizabeth Tseng, Preston Tuggle, Nick Turley, Jerry Tworek, Juan Felipe Cerón Uribe, Andrea Vallone, Arun Vijayarvigiya, Chelsea Voss, Carroll Wainwright, Justin Jay Wang, Alvin Wang, Ben Wang, Jonathan Ward, Jason Wei, CJ Weinmann, Akila Welihinda, Peter Welinder, Jiayi Weng, Lilian Weng, Matt Wiethoff, Dave Willner, Clemens Winter, Samuel Wolrich, Hannah Wong, Lauren Workman, Sherwin Wu, Jeff Wu, Michael Wu, Kai Xiao, Tao Xu, Sarah Yoo, Kevin Yu, Qiming Yuan, Wojciech Zaremba, Rowan Zellers, Chong Zhang, Marvin Zhang, Shengjia Zhao, Tianhao Zheng, Juntang Zhuang, William Zhuk, and Barret Zoph. 2024. GPT-4 Technical Report. arXiv:2303.08774 [cs.CL]
- [41] Oracle. 2020. *Deprecated List (Java SE 15 & JDK 15)*. <https://docs.oracle.com/en/java/javase/15/docs/api/deprecated-list.html>
- [42] Peter-Michael Osera and Steve Zdancewic. 2015. Type-and-example-directed program synthesis. In *Proceedings of the 36th ACM SIGPLAN Conference on Programming Language Design and Implementation, Portland, OR, USA, June 15-17, 2015*, David Grove and Stephen M. Blackburn (Eds.). ACM, 619–630. doi:10.1145/2737924.2738007
- [43] Rohan Padhye, Caroline Lemieux, and Koushik Sen. 2019. JQF: coverage-guided property-based testing in Java. In *Proceedings of the 28th ACM SIGSOFT International Symposium on Software Testing and Analysis, ISSTA 2019, Beijing, China, July 15-19, 2019*, Dongmei Zhang and Anders Möller (Eds.). ACM, 398–401. doi:10.1145/3293882.3339002
- [44] Rangeet Pan, Ali Reza Ibrahimzade, Rahul Krishna, Divya Sankar, Lambert Pouguem Wassi, Michele Merler, Boris Sobolev, Raju Pavuluri, Saurabh Sinha, and Reyhaneh Jabbarvand. 2024. Lost in Translation: A Study of Bugs Introduced by Large Language Models while Translating Code. In *Proceedings of the IEEE/ACM 46th International Conference on Software Engineering (Lisbon, Portugal), (ICSE '24)*. Association for Computing Machinery, New York, NY, USA, Article 82, 13 pages. doi:10.1145/3597503.3639226
- [45] Jeff H. Perkins. 2005. Automatically generating refactorings to support API evolution. In *Proceedings of the 2005 ACM SIGPLAN-SIGSOFT Workshop on Program Analysis For Software Tools and Engineering, PASTE'05, Lisbon, Portugal, September 5-6, 2005*, Michael D. Ernst and Thomas P. Jensen (Eds.). ACM, 111–114. doi:10.1145/1108792.1108818
- [46] Baptiste Rozière, Jonas Gehring, Fabian Gloeckle, Sten Sootla, Itai Gat, Xiaoqing Ellen Tan, Yossi Adi, Jingyu Liu, Romain Sauvestre, Tal Remez, Jérémy Rapin, Artyom Kozhevnikov, Ivan Evtimov, Joanna Bitton, Manish Bhatt, Cristian Canton Ferrer, Aaron Grattafiori, Wenhan Xiong, Alexandre Défossez, Jade Copet, Faisal Azhar, Hugo Touvron, Louis Martin, Nicolas Usunier, Thomas Scialom, and Gabriel Synnaeve. 2024. Code Llama: Open Foundation Models for Code. arXiv:2308.12950 [cs.CL]
- [47] Max Schäfer, Julian Dolby, Manu Sridharan, Emina Torlak, and Frank Tip. 2010. Correct Refactoring of Concurrent Java Code. In *ECOOP 2010 – Object-Oriented Programming*, Theo D'Hondt (Ed.). Springer Berlin Heidelberg, Berlin, Heidelberg, 225–249.
- [48] Armando Solar-Lezama, Christopher Grant Jones, and Rastislav Bodík. 2008. Sketching concurrent data structures. In *Proceedings of the ACM SIGPLAN 2008 Conference on Programming Language Design and Implementation, Tucson, AZ, USA, June 7-13, 2008*, Rajiv Gupta and Saman P. Amarasinghe (Eds.). ACM, 136–148. doi:10.1145/1375581.1375599
- [49] Friedrich Steimann. 2011. Constraint-Based Model Refactoring. In *Model Driven Engineering Languages and Systems: 14th International Conference (MODELS)*, Jon Whittle, Tony Clark, and Thomas Kühne (Eds.). Springer, 440–454. doi:10.1007/978-3-642-24485-8_32
- [50] Friedrich Steimann, Christian Kollee, and Jens von Pilgrim. 2011. A Refactoring Constraint Language and Its Application to Eiffel. In *ECOOP 2011 – Object-Oriented Programming: 25th European Conference*, Mira Mezini (Ed.). Springer, 255–280. doi:10.1007/978-3-642-22655-7_13
- [51] Friedrich Steimann and Jens von Pilgrim. 2012. Constraint-Based Refactoring with Foresight. In *ECOOP 2012 – Object-Oriented Programming: 26th European Conference*, James Noble (Ed.). Springer, 535–559. doi:10.1007/978-3-642-31057-7_24
- [52] Zilu Tang, Mayank Agarwal, Alexander Shypula, Bailin Wang, Derry Wijaya, Jie Chen, and Yoon Kim. 2023. Explain-then-translate: an analysis on improving program translation with self-generated explanations. In *Findings of the Association for Computational Linguistics: EMNLP 2023*, Houda Bouamor, Juan Pino, and Kalika Bali (Eds.). Association for Computational Linguistics, Singapore, 1741–1788. doi:10.18653/v1/2023.findings-emnlp.119
- [53] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmin Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shrutit Bhosale, Dan Bikel, Lukas Blecher, Cristian Canton Ferrer, Moya Chen, Guillem Cucu-rull, David Esiobu, Jude Fernandes, Jeremy Fu, Wenyin Fu, Brian Fuller, Cynthia Gao, Vedanuj Goswami, Naman Goyal, Anthony Hartshorn, Saghar Hosseini, Rui Hou, Hakan Inan, Marcin Kardas, Viktor Kerkez, Madsen Khabsa, Isabel Kloumann, Artem Korenev, Punit Singh Koura, Marie-Anne Lachaux, Thibaut Lavril, Jenya Lee, Diana Liskovich, Yinghai Lu, Yuning Mao, Xavier Martinet, Todor Mihaylov, Pushkar Mishra, Igor Molybog, Yixin Nie, Andrew Poulton, Jeremy Reizenstein, Rashi Rungta, Kalyan Saladi, Alan Schelten, Ruan Silva, Eric Michael Smith, Ranjan Subramanian, Xiaoqing Ellen Tan, Binh Tang, Ross Taylor, Adina Williams, Jian Xiang Kuan, Puxin Xu, Zheng Yan, Iliyan Zarov, Yuchen Zhang, Angela Fan, Melanie Kambadur, Sharan Narang, Aurelien Rodriguez, Robert Stojnic, Sergey Edunov, and Thomas Scialom. 2023. Llama 2: Open Foundation and Fine-Tuned Chat Models. arXiv:2307.09288 [cs.CL]
- [54] Elco Visser. 2004. *Program Transformation with Stratego/XT. Rules, Strategies, Tools, and Systems in Stratego/XT 0.9*. Technical Report UU-CS-2004-011. Department of Information and Computing Sciences, Utrecht University.
- [55] Wei Wu, Yann-Gaël Guéhéneuc, Giuliano Antoniol, and Miryung Kim. 2010. AURA: a hybrid approach to identify framework evolution. In *Proceedings of the 32nd ACM/IEEE International Conference on Software Engineering – Volume 1, ICSE 2010, Cape Town, South Africa, 1-8 May 2010*, Jeff Kramer, Judith Bishop, Premkumar T. Devanbu, and Sebastián Uchitel (Eds.). ACM, 325–334. doi:10.1145/1806799.1806848
- [56] Frank F. Xu, Uri Alon, Graham Neubig, and Vincent J. Hellendoorn. 2022. A Systematic Evaluation of Large Language Models of Code. CoRR abs/2202.13169 (2022). arXiv:2202.13169 <https://arxiv.org/abs/2202.13169>
- [57] Masao Yamaguchi, Kazutaka Matsuda, Cristina David, and Meng Wang. 2021. Synbit: synthesizing bidirectional programs using unidirectional sketches. *Proc. ACM Program. Lang.* 5, OOPSLA (2021), 1–31. doi:10.1145/3485482
- [58] John Yang, Carlos E. Jimenez, Alexander Wettig, Kilian Lieret, Shunyu Yao, Karthik Narasimhan, and Ofir Press. 2024. SWE-agent: Agent-Computer Interfaces Enable Automated Software Engineering. arXiv:2405.15793 [cs.SE]
- [59] Fengji Zhang, Bei Chen, Yue Zhang, Jacky Keung, Jin Liu, Daoguang Zan, Yi Mao, Jian-Guang Lou, and Weizhu Chen. 2023. RepoCoder: Repository-Level Code Completion Through Iterative Retrieval and Generation. In *The 2023 Conference on Empirical Methods in Natural Language Processing*. <https://openreview.net/forum?id=q09vTY1Cqh>
- [60] Hanliang Zhang, Cristina David, Wang Meng, Brandon Paulsen, and Daniel Kroening. 2025. Scalable, Validated Code Translation of Entire Projects using Large Language Models. *Programming Language Design and Implementation (PLDI) (2025)*.
- [61] Hanliang Zhang, Cristina David, Yijun Yu, and Meng Wang. 2023. Ownership Guided C to Rust Translation. In *Computer Aided Verification - 35th International Conference, CAV 2023, Paris, France, July 17-22, 2023, Proceedings, Part III (Lecture Notes in Computer Science, Vol. 13966)*, Constantin Enea and Akash Lal (Eds.). Springer, 459–482. doi:10.1007/978-3-031-37709-9_22
- [62] Shuyan Zhou, Uri Alon, Sumit Agarwal, and Graham Neubig. 2023. CodeBERTScore: Evaluating Code Generation with Pretrained Models of Code. CoRR abs/2302.05527 (2023). doi:10.48550/arXiv.2302.05527 arXiv:2302.05527
- [63] Qihao Zhu, Qingyuan Liang, Zeyu Sun, Yingfei Xiong, Lu Zhang, and Shengyu Cheng. 2024. GrammarT5: Grammar-Integrated Pretrained Encoder-Decoder Neural Model for Code. In *Proceedings of the IEEE/ACM 46th International Conference on Software Engineering (Lisbon, Portugal), (ICSE '24)*. Association for Computing Machinery, New York, NY, USA, Article 76, 13 pages. doi:10.1145/3597503.3639125