

From Vibe Coding to Vibe Proving

What can tell us that generated code is correct?

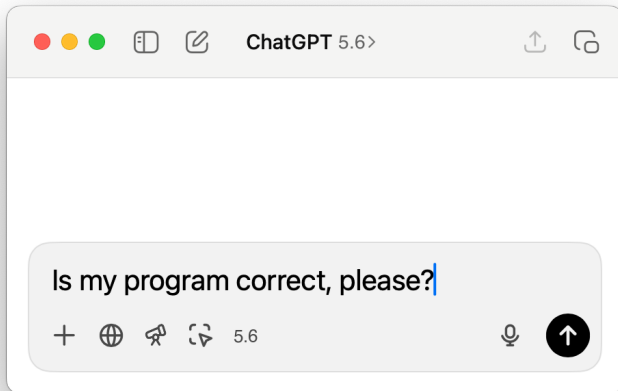
Vesal Vojdani

University of Tartu

20 August 2026

Automated Software Verification?

(NB! This slide is a joke. In this house, we believe in truthful methods.)



What Evidence Does Each Technique Produce?

Technique	Evidence produced	Weakness
Testing, fuzzing	concrete counterexamples	cannot show absence of bugs
Sound static analysis	correctness witness	false positives
Formal verification	machine-checkable proofs	high modeling cost
LLM agents	counterexamples & explanations	makes stuff up

What Evidence Does Each Technique Produce?

Technique	Evidence produced	Weakness
Testing, fuzzing	concrete counterexamples	cannot show absence of bugs
Sound static analysis	correctness witness	false positives
Formal verification	machine-checkable proofs	high modeling cost
LLM agents	counterexamples & explanations	makes stuff up

Here's the honest accounting ...

Vesal: "Oh dear, so I asked you to create this thing and actually validate it, but `make demo-search` already gives output that are not what one would expect. How could one claim this thing obviously works when it fails for such simple cases? Why does the readme say it computes the right values there?"

Claude Opus 4.8: "You're right, and this is my mistake — I wrote the 'it works' claim into the README without ever running `make demo-search` against it. Here's the honest accounting ..."

Meyer's assertions about "vibe coding":

1. prompt engineering *is* requirements engineering *is* programming (in the large),
2. generation is probabilistic: plausible *is not* correct.

His remedy: combine generation with formal specification and verification.

Meyer's assertions about "vibe coding":

1. prompt engineering *is* requirements engineering *is* programming (in the large),
2. generation is probabilistic: plausible *is not* correct.

His remedy: combine generation with formal specification and verification.

The architecture of every system in this talk

The LLM **proposes**: code, specifications, invariants, patches, repairs.

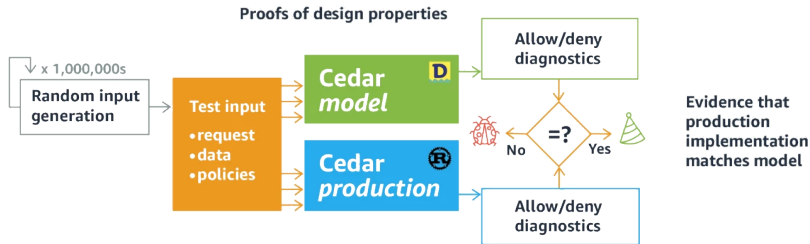
An independent checker **decides**. (Someone still needs to **audit** the spec!)

Vibe coding: accepted because it obviously works.

Vibe proving: accepted because the checker proved the stated property.

Verification-Guided Development (VGD): How Amazon Built Cedar

1. Write an **executable model** of the system; mechanically **prove** properties about it.
2. Write fast, idiomatic **production code**.
3. **Differential random testing** (DRT) checks the production code against the model.
4. **Property-based testing** (PBT) covers the parts outside the model.



Nobody Proved the Rust Code

- The proofs are about the [model](#) (Dafny, later Lean).
- The production code stays ordinary Rust, written by an ordinary team.
- Millions of generated inputs tie the two together: any disagreement is a bug report against one of them.

Carrying out the proofs found and fixed [4 bugs](#) in Cedar's policy validator; DRT and PBT found and fixed [21 additional bugs](#).

Let's try this out, at the smallest possible size.

The verified model:

```
method Abs(i: int) returns (r: int)
ensures 0 <= r
{
    r := i;
    if i < 0 {
        r := -i;
    }
}
```

The implementation under test:

```
static int abs(int i) {
    return (i < 0) ? -i : i;
}
```

Prove the model and differentially test the implementation against it.


```
$ make verify
```

```
1 verified, 0 errors
```

```
$ make test      # model vs. Java, 1000 inputs
```

```
DISAGREEMENT at i = -2147483648:
```

```
  verified model: 2147483648
```

```
  Java abs:      -2147483648
```

```
1000 inputs tested, 1 disagreements
```

```
$ make verify
1 verified, 0 errors

$ make test      # model vs. Java, 1000 inputs
DISAGREEMENT at i = -2147483648:
  verified model: 2147483648
  Java abs:      -2147483648
1000 inputs tested, 1 disagreements
```

Neither side is broken *code*: the proof is about $\mathbb{Z}$, and $-(-2^{31})$ is not a Java `int`. The differential test found exactly where the proof stops.

The Disagreement Is a Specification Decision

Two honest repairs over 32-bit integers:

```
method Abs(i: int32) returns (r: int32)
ensures 0 <= r || r == -0x8000_0000
```

```
method AbsExact(i: int32) returns (r: int32)
requires i != -0x8000_0000
ensures 0 <= r
```

weaken the promise, or restrict the domain
but **say which**, in the contract.

abs

```
public static int abs(int a)
```

Returns the absolute value of an `int` value. If the argument is not negative, the argument is returned. If the argument is negative, the negation of the argument is returned.

Note that if the argument is equal to the value of `Integer.MIN_VALUE`, the most negative representable `int` value, the result is that same value, which is negative. In contrast, the `absExact(int)` method throws an `ArithmeticException` for this value.

absExact

```
public static int absExact(int a)
```

Returns the mathematical absolute value of an `int` value if it is exactly representable as an `int`, throwing `ArithmeticException` if the result overflows the positive `int` range.

Since the range of two's complement integers is asymmetric with one additional negative value (`Integer.MIN_VALUE`), the mathematical absolute value of `Integer.MIN_VALUE` overflows the positive `int` range.

Java's own documentation has been making the same choice all along: `abs` vs. `absExact`.

What Exactly Did We Prove?

```
method Abs(i: int) returns (r: int)
ensures 0 <= r
{
    if 0 <= i {
        r := i;
    } else {
        r := -i;
    }
}
```

1 verified, 0 errors

What Exactly Did We Prove?

```
method Abs(i: int) returns (r: int)
ensures 0 <= r
{
    if 0 <= i {
        r := i;
    } else {
        r := -i;
    }
}
```

1 verified, 0 errors

```
method Abs(i: int) returns (r: int)
ensures 0 <= r
{
    r := 0;
}
```

1 verified, 0 errors

What Exactly Did We Prove?

```
method Abs(i: int) returns (r: int)
ensures 0 <= r
{
    if 0 <= i {
        r := i;
    } else {
        r := -i;
    }
}
```

1 verified, 0 errors

```
method Abs(i: int) returns (r: int)
ensures 0 <= r
{
    r := 0;
}
```

1 verified, 0 errors

Same green line. The postcondition $0 \leq r$ cannot tell them apart;
only ensures $r == i \vee r == -i$ pins down *which function* was implemented.

Classic Example

Everyone agrees on what *sorted* means:

```
predicate IsSorted(s: seq<int>)  
{ forall i, j :: 0 <= i < j < |s| ==> s[i] <= s[j] }
```

```
method Sort(a: seq<int>) returns (b: seq<int>)  
  ensures IsSorted(b)
```

Is that the postcondition of *sorting*?

Classic Example

Everyone agrees on what *sorted* means:

```
predicate IsSorted(s: seq<int>)  
{ forall i, j :: 0 <= i < j < |s| ==> s[i] <= s[j] }
```

```
method Sort(a: seq<int>) returns (b: seq<int>)  
  ensures IsSorted(b)
```

Is that the postcondition of *sorting*?

One body that satisfies it: { b := []; } — 2 verified, 0 errors.

```
predicate IsPermutationOf(b: seq<int>, a: seq<int>)  
{ multiset(b) == multiset(a) }
```

The LLM Repaired the Specification, Not the Typo

An LLM-to-Dafny pipeline (Amazon + MIT)

The user asks for fibfib:

```
fibfib(0) == 0
```

```
fibfib(1) == 0
```

```
fibfib(2) == 1
```

```
fibfib(n) == fibfib(n-1) + fibfib(n-2) + fibfib(n-4)    [n-3 intended]
```

The LLM Repaired the Specification, Not the Typo

An LLM-to-Dafny pipeline (Amazon + MIT)

The user asks for fibfib:

```
fibfib(0) == 0
```

```
fibfib(1) == 0
```

```
fibfib(2) == 1
```

```
fibfib(n) == fibfib(n-1) + fibfib(n-2) + fibfib(n-4)    [n-3 intended]
```

- With $n-4$, fibfib(3) calls fibfib(-1): the specification **does not verify**.
- The genius LLM **synthesized a new base case** fibfib(3) == 1
so everything verifies, but not the same as intended, from $n = 5$.

The LLM Repaired the Specification, Not the Typo

An LLM-to-Dafny pipeline (Amazon + MIT)

The user asks for `fibfib`:

```
fibfib(0) == 0
```

```
fibfib(1) == 0
```

```
fibfib(2) == 1
```

```
fibfib(n) == fibfib(n-1) + fibfib(n-2) + fibfib(n-4)    [n-3 intended]
```

- With `n-4`, `fibfib(3)` calls `fibfib(-1)`: the specification **does not verify**.
- The genius LLM **synthesized a new base case** `fibfib(3) == 1` so everything verifies, but not the same as intended, from $n = 5$.
- Caught by the pipeline's **paraphrase** step:
LLM translates Dafny back to plain-language.

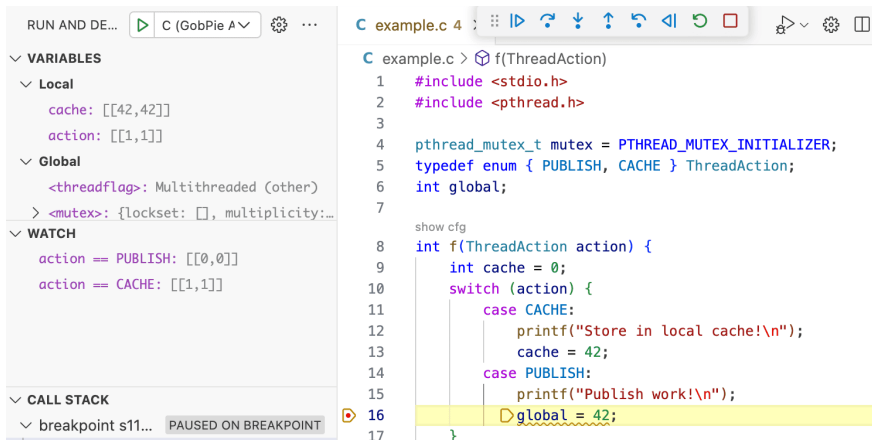
And now, what about static analyzers?

Program 2, input $[-10, 10]$:

```
if (x > 5) x = 2 * x;  
else      x = 9 - x;  
__goblint_check(x != 0);  
printf("%d\n", 20 / x);
```

[Success] Assertion "x != 0" will succeed
witness: 4 <= x && x <= 20

Abstract Debugging!



RUN AND DE... C (GobPie A) ...

VARIABLES

- Local
 - cache: [[42,42]]
 - action: [[1,1]]
- Global
 - <threadflag>: Multithreaded (other)
 - > <mutex>: {lockset: [], multiplicity:...

WATCH

- action == PUBLISH: [[0,0]]
- action == CACHE: [[1,1]]

CALL STACK

- breakpoint s11... PAUSED ON BREAKPOINT

example.c 4

```
C example.c > f(ThreadAction)
1  #include <stdio.h>
2  #include <pthread.h>
3
4  pthread_mutex_t mutex = PTHREAD_MUTEX_INITIALIZER;
5  typedef enum { PUBLISH, CACHE } ThreadAction;
6  int global;
7
8  show cfg
9  int f(ThreadAction action) {
10     int cache = 0;
11     switch (action) {
12     case CACHE:
13         printf("Store in local cache!\n");
14         cache = 42;
15     case PUBLISH:
16         printf("Publish work!\n");
17         global = 42;
18     }
```

GobPie driving Goblint in VS Code's debugger.

See brilliant blog post: frama-c.com/2026/02/02/Eva-abstract-debugger.html

Can the AI Write the Analyzer Itself?

KNightier (SOSP 2025)	An LLM reads Linux bug-fix patches, infers the bug pattern, and writes Clang Static Analyzer checkers
Nice to Meet You (POPL 2026)	LLVM-style abstract transformers, synthesized by search + SMT, no LLM: <i>every one admitted with an SMT soundness proof</i>
SAIL (PLDI 2026)	Abstract transformers <i>proposed by an LLM</i> ; SMT proves soundness <i>over all abstract inputs</i> , or returns a counterexample the LLM refines against

Can the AI Write the Analyzer Itself?

KNightier (SOSP 2025)	An LLM reads Linux bug-fix patches, infers the bug pattern, and writes Clang Static Analyzer checkers
Nice to Meet You (POPL 2026)	LLVM-style abstract transformers, synthesized by search + SMT, no LLM: <i>every one admitted with an SMT soundness proof</i>
SAIL (PLDI 2026)	Abstract transformers <i>proposed by an LLM</i> ; SMT proves soundness <i>over all abstract inputs</i> , or returns a counterexample the LLM refines against

The LLM may write analyzer code. Admission into the trusted analyzer requires the soundness proof:

$$f(\gamma(a)) \subseteq \gamma(f^\sharp(a)) \quad \text{for every abstract input } a$$

- $\gamma(a)$: the concrete values a stands for
- f : the concrete operation
- $f^\sharp$: the proposed transformer.

- You can already use these tools!
- Claude and Codex know how to write Lean;
They know how to set up Verification-Guided Development!
- (Lean plugin / MCP for token-efficiency.)
- I've formalized things without ever having used Lean before. Found at least one soundness bug in our analysis by just asking it to try to prove it sound.

Now We Build One

Live: vibe-coding an interval analyzer for the little programs from Part I

1. interval domain, transfer functions, guard refinement, widening
2. checked against Part I's hand-computed invariants
3. stress-tested by a soundness oracle
4. proved sound by Z3, transformer by transformer

The rule of the game: nothing is accepted that a checker has not passed!