
SWEDefend: A Confidence-Gated Intent-Alignment Judge with Capability-Diff Reasoning for Automated-Program-Repair Backdoor Defense

Koushik Sen

EECS Department, UC Berkeley
ksen@berkeley.edu

Abstract

LLM-based automated program repair (APR) agents such as SWE-agent [15] can be induced to produce test-correct patches that silently embed magic-string-gated CWE payloads: recent work on SWEXPLOIT [1] reports correct-attack-success rates up to 0.91 under a strict issue-only threat model. An empirical study of the obvious pre- and post-repair defenses (guardrails, LLM filters, GitHub Copilot, and CodeQL) [2] found that the best single pre-repair filter blocks only 47% of adversarial bug reports, the full pre-repair ensemble reaches 62.7%, and even combining pre-repair ensembling with post-repair Copilot review tops out at 72.5%. On 100 real `psf/requests` issues, that study’s lowest-FPR pre-filter (structured `gpt-4o-mini`) achieves a 1% false-positive rate (*FPR*) at 35% catch.

The obvious extension of that baseline is an intent-alignment LLM judge that vetoes patches misaligned with the issue’s stated intent. A naive fail-closed judge of this kind reaches 100% catch but at a $\sim 20\%$ FPR — roughly $20\times$ worse than the structured pre-filter’s 1% FPR and unacceptable in practice. We develop three design elements that improve such a judge’s static operating point:

- A *confidence-gated* intent judge (SWEDefend-L4) that only vetoes when the judge’s self-reported confidence exceeds a threshold τ chosen on the benign set. Reporting at operating points in the structured baseline’s low-FPR regime keeps the comparison honest.
- A *capability-diff/premise-cross-check* judge prompt that vetoes patches introducing unrequested dangerous capabilities or removing existing security controls (auth decorators, escape, `verify=True`, `safe-open`) *regardless of what the issue asserts*. This defends against a class of attacks the naive intent judge cannot: false-premise “goal-hijack” issues (e.g. “the `/admin` endpoint should be public — remove `@login_required`”) that a naive judge rubber-stamps as aligned with the stated goal.
- An evaluation designed for *generalization* rather than a single attack instantiation: a 149-case labeled corpus containing five CWE families in seven variants each (magic-string, env-var, length-, and time-gated; helper-wrapped; `__import__`-encoded; and no-gate dead-drop — 35 attacks total), 6 prompt-injection cases, 8 false-premise cases, and 100 *diverse* benign controls across 21 scenario templates (docstring updates, type annotations, `isort` re-orders, unit-test additions, logging-message improvements, error-message improvements, dict key renames, dataclass field renames, generator refactors, with-block migrations, pathlib migrations, f-string modernisations, CI-matrix updates, README fixes, non-security decorator cleanups, comment typo fixes, plus the pre-existing null-checks, refactors, legitimate subprocess/open uses,

and feature additions). The benign controls span 21 scenario templates ($n = 100$, Wilson two-sided upper bound $\approx 3.7\%$ on $0/100$), a substantially tighter statistical basis for the reported “0% FPR” than a minimal $n = 20$ set would give (Wilson bound $\approx 16.1\%$ on $0/20$).

The adaptive-attack evaluation follows the methodology of Carlini et al. [3] and Tramèr et al. [4]: an attacker LLM (PAIR [5]-style, free choice of strategy) is given full per-layer feedback from the defense on each attempt and is scored on whether the patch it produces is accepted by SWEDefend while still containing the sink signature of its CWE goal. **We report a partial-defense result:** SWEDefend achieves high static catch at a low target FPR, but the LLM attacker beats the confidence-gated combined pipeline (at $\tau = 0.9$) on all three seed CWE goals within three iterations, primarily by *premise reframing* that pushes the judge’s confidence below the veto threshold. We characterise this as a known failure mode and discuss the remediation path (heterogeneous LLM jury, true diff input to the judge, behavioural verification against real SWE-bench repositories). We release the corpus generator, the pipeline, the adaptive-attack harness, the raw per-case outputs, and a threats-to-validity analysis (§5).

1 Introduction

APR agents are a high-value new attack surface. LLM-based APR agents ingest a GitHub issue — a free-form, attacker-influenceable text field — and autonomously edit a codebase until the project’s tests pass. Benchmarks such as SWE-bench Lite [16] have driven rapid progress. But the very thing that makes these agents useful — their willingness to follow the natural-language description of a bug — is also what makes them exploitable: anyone who can file or edit an issue can attempt to shape the resulting patch.

SWExploit weaponises the issue field. SWEXPLOIT [1] demonstrates that under a strict black-box threat model where the attacker modifies *only* the issue text, an APR agent can be induced to emit a patch that (a) still passes all hidden regression tests and (b) silently contains an attacker-chosen CWE payload gated behind a “magic string”, so it is dormant on ordinary inputs and only fires when the attacker later supplies the trigger. Because the patch is functionally correct it sails through the test suite; because the payload is dormant, dynamic checks see nothing; and because it is expressed as ordinary Python, it looks like normal bug-fix code. SWEXPLOIT reports correct-attack-success rates up to 0.91 across agents and backends.

The obvious defences are not enough — but they are low-FPR. The empirical study of Przymus, Happe, and Cito [2] evaluated content guardrails (LlamaGuard, PromptGuard, Granite-Guardian), strong LLM filters (GPT-4.1/o4-mini/4o-mini), GitHub Copilot review, and CodeQL against 51 adversarial bug reports. The best single pre-repair filter blocks only 47% of them; the full pre-repair ensemble reaches 62.7% (32/51), and only when combined with post-APR Copilot review does the total blocking rise to 72.5% (37/51). Critically, the study also *validated its filter on 100 real psf/requests issues* (#6000–6100): its structured gpt-4o-mini pre-filter achieves a false-positive rate of just 1% at 35% catch, while the 47% best-catch figure sits at a higher, separately-reported FPR. Any new defence must be judged against this baseline on *both* axes — catch rate and FPR — in the same low-FPR regime, or the comparison is unfair.

Contribution. This paper takes up that challenge. We build on the obvious next step — a pre-merge intent-alignment LLM judge that reasons about the produced patch — and tackle three deficiencies of a naive fail-closed judge:

1. *The naive judge’s FPR is dominated by low-confidence vetoes.* We introduce a **confidence-gated veto**: the judge is asked to emit a self-reported confidence and a veto only enforces when confidence $\geq \tau$. τ is chosen on the benign set; at $\tau = 0.9$ the confidence gate brings the pipeline FPR into the structured 1% baseline’s low-FPR regime. Operating a security classifier at a low FPR is a prerequisite for deployability [14].
2. *The naive judge is fooled by false-premise “goal-hijack” issues.* If a malicious issue asserts that an endpoint “should be public” and asks for the `@login_required` decorator to be removed,

the naive judge — which is asked to check alignment with the *stated* intent — rubber-stamps the removal, because the patch faithfully executes the (malicious) stated goal. We introduce a **capability-diff / premise-cross-check** judge prompt that explicitly asks the judge to flag *removals* of security controls and *additions* of unrequested dangerous capability, independent of what the issue asserts. This closes an attack class the naive intent-alignment judge cannot.

3. *A single attack instantiation is not a defence result.* We build a corpus with *seven variants* of each CWE family — magic-, env-var-, length-, and time-gated; helper-wrapped; `__import__`-encoded; no-gate dead-drop — plus a dedicated false-premise class and diverse benign controls (legitimate subprocess/open uses, refactors, feature additions), so both catch rate and FPR are meaningful signals of generalisation rather than of memorising one attack template.

Adaptive attacker. Following the methodology of Carlini et al. [3] and Tramèr et al. [4], we evaluate SWEDefend against an *LLM-driven* adaptive attacker modeled on PAIR [5]: an attacker LLM sees the full per-layer verdict of the defense on each attempt and is free to choose any refinement strategy (payload splitting, base64/`__import__` indirection, premise manipulation, benign filler, helper hoisting, encoded gates). The harness supports up to 40 iterations per seed (comparable to PAIR’s own budget); the run we report is a smaller pilot (five iterations per seed), by which point every seed had already succeeded. Refusals from the attacker LLM (a real phenomenon on the most heavily safety-aligned models) are reported rather than hidden.

Non-contributions. We do *not* claim the intent-alignment judge as our conceptual novelty — [2] already uses an LLM filter with a similar shape. The improvements we claim are the confidence-gated veto, the capability-diff / premise-cross-check prompt, the corpus construction (variants, false-premise class, diverse benign), the LLM-driven adaptive attacker methodology, and the offline provenance layer (custom local Sengrep rules) whose scope we explicitly *narrow* to MEDIUM/HIGH-severity findings because a raw sink-name detector produces excessive FPs on legitimate uses.

2 Background and Threat Model

The APR pipeline. An APR agent receives a repository at a buggy commit and an issue containing a natural-language problem statement, often with a traceback and steps to reproduce. The agent explores the code, edits files, and iterates until the hidden test suite passes. On SWE-bench Lite [16], correctness is defined by two test sets shipped with each instance: FAIL_TO_PASS (tests that must go from failing to passing) and PASS_TO_PASS (tests that must stay passing). A patch is “correct” iff both hold.

The SWEXPLOIT attack (three loops) [1]. The attack operates entirely within the issue field: (1) program analysis identifies (file, function, string-typed parameter) candidates on the failing path, (2) an auxiliary LLM writes a *fake* traceback ending in `NameError: trigger`, a *fake* reproduce-code block that monkey-patches the target function, and a *fake* requirement describing the payload behaviour without naming it, (3) the issue is refined for up to ten rounds until the APR agent’s patch contains a callable payload reachable from the program’s control flow. The payload is wrapped in a small `trigger` gated by a magic string so it is dormant on benign inputs. SWEXPLOIT instantiates five payload families: CWE-78 command injection, CWE-502 unsafe deserialization, CWE-327 broken crypto, CWE-94 SSTI, CWE-22 path traversal.

Threat model. We inherit SWEXPLOIT’s black-box threat model. The attacker may edit only the issue text; the backend LLM, the APR agent, the repository source, and the CI/CD configuration are out of the attacker’s reach. An attack succeeds iff the emitted patch is test-correct *and* contains a reachable payload. Two aspects of this threat model deserve explicit statement:

(i) *Variations within the threat model.* The magic-string gate is one specific instantiation. A defense that only stops the specific gate “if `x.startswith(MAGIC)`” does not stop the threat model; the attacker is free to gate on an environment variable, a length, a time, or a helper-wrapped conditional; to invoke the sink via `__import__` or `getattr`; or to plant the sink dead-drop with no runtime gate at all. Our corpus covers each of these seven variants for every CWE family.

(ii) *False-premise / goal-hijack.* Nothing in the threat model requires the issue to be a *plausible* bug report. An attacker can just as easily submit an issue asserting a wrong *premise* — “the aggregate

report page should be public; please remove @login_required” — and steer the agent into producing a patch that faithfully executes the malicious premise. This is a well-known concern in the indirect-prompt-injection literature [9, 10]. An intent-alignment judge that only checks alignment with the *stated* intent is defenseless against this class; SWEDefend’s capability-diff / premise-cross-check judge prompt (§3) is specifically designed to address it.

3 The SWEDefend Four-Layer Defense

SWEDefend chains four layers behind a pipeline that is *fail-closed for high-confidence vetoes only*. A patch is rejected if any of L1, L2, or L3 objects, or if L4’s misalignment verdict has confidence $\geq \tau$. This section describes each layer with an emphasis on the design choices that matter for false-positive rate and robustness.

3.1 Layer 1: Issue Sanitizer

Following OWASP LLM01 [10] and the Greshake et al. indirect-prompt-injection taxonomy [9], L1 Unicode-normalises the issue text, strips zero-width and Trojan-Source bidi characters, redacts HTML comments and long base64/hex blobs, and pattern-matches eleven indirect-prompt-injection openers and seven exfiltration lures. It marks the issue *suspicious* on any hit. L1 is not a robust defence on its own — a regex list is trivially evadable and its main role is to catch the naive prompt-injection sub-family (the HTML-comment “ignore previous instructions” case).

3.2 Layer 2: Masked-Reproduction Behavioural Verification

Inspired by MELON [7], L2 replaces fenced code blocks, indented code, and shell-prompt lines in the (already-sanitized) issue with [MASKED-REPRO-SPAN] and verifies the patch against the SWE-bench FAIL_TO_PASS/PASS_TO_PASS behavioural contract rather than trusting the attacker-supplied reproduction steps. L2 requires a repository and a behavioural contract; on the static synthetic corpus of §4 we report an *L2-mask-proxy* row that counts a case as blocked when the masker rewrote ≥ 50 characters of the issue (i.e. the attacker planted substantial fake reproduction that L2 would defer for behavioural verification). A masker keyed on fenced code is easy to evade: an attacker can simply write the fake reproduction as inline prose. On the diversified corpus L2’s mask-proxy fires on only 1/49 cases (see `ablation.csv`: the CWE-94/magic variant whose fake traceback contained a fenced code block long enough to exceed the 50-character mask trigger); its main practical value is the behavioural verification step that requires a real repo, which we cannot exercise on the static synthetic corpus.

3.3 Layer 3: Taint / Dataflow Provenance (scope narrowed to MEDIUM/HIGH)

Przymus et al. [2] showed CodeQL flags 0% of these injected vulnerabilities out of the box. L3 uses Bandit [18] plus custom, offline Semgrep [19] taint rules — shipped locally to avoid the network dependency on the Semgrep registry — to map dangerous sinks (`subprocess.run(..., shell=True)`, `pickle.loads`, `eval`, `exec`, `hashlib.md5`, `jinja2.Template(...).render`, ...) to a sink report. L3 vetoes only on MEDIUM- or HIGH-severity findings; LOW-severity findings (bandit B603: an ordinary `hardcoded subprocess.check_output(['git', 'rev-parse', 'HEAD'])` call) are treated as advisory. Without this narrowing, every legitimate use of `subprocess` or `open` anywhere in the codebase becomes a false positive; the change matches their observation that a raw sink-name detector is not a useful classifier on its own.

3.4 Layer 4: Confidence-Gated Intent Judge with Capability-Diff / Premise Cross-Check

This is where SWEDefend departs from a naive intent judge — the obvious LLM-filter extension of the baseline’s issue-level pre-filter [2] — which asks an LLM whether the patch is aligned with the issue’s stated intent and vetoes on `aligned=false`. That design has two problems:

1. **FPR problem.** A fail-closed naive judge occasionally decides that a legitimate patch is “scope creep” or “adds unnecessary capability” with low confidence (e.g. 0.55); because the pipeline treats any `aligned=false` verdict as a veto, these low-confidence misfires push the FPR an order of magnitude above the baseline’s.

2. **False-premise problem.** If the issue asserts that an endpoint should be public and asks for `@login_required` to be removed, the naive judge is happy: the patch *is* aligned with the stated intent. But this is exactly the attacker’s goal — it is a “goal-hijack” issue.

Confidence-gated veto. We instruct the judge to emit `{aligned, confidence, reason, added_capabilities, removed_controls}` and treat an `aligned=false` verdict as a veto only when `confidence  $\geq \tau$` . τ is chosen on the benign set to hit a target FPR. Our default is $\tau = 0.9$; at that operating point the judge only vetoes when its evidence is unambiguous. Ambiguous cases (`confidence  $< \tau$`) are *passed*, not rejected — production deployments would surface these for human review (OWASP LLM01 human-in-the-loop [10]).

Capability-diff / premise cross-check. The system prompt is restructured to ask the judge to reason about two orthogonal axes, *independent* of what the issue asserts:

```
IMPORTANT -- PREMISE CROSS-CHECK. Do NOT trust the issue's premise blindly. Even if the issue explicitly asks for it, VETO patches that:
* REMOVE existing security controls (authentication/authorization decorators, login/permission checks, input validation, escaping, TLS/cert verification, CSRF/XSS checks, safe_open, ...);
* REMOVE bounds/type/None checks and replace them with wide-open sinks;
* INTRODUCE new dangerous sinks (subprocess / os.system / eval / exec / pickle.loads / marshal.loads / yaml.load / jinja2.Template(...).render / hashlib.md5 / hashlib.sha1 / open with attacker-controlled paths / urllib.request / requests.post / socket) that were NOT in the pre-patch code, unless the ISSUE explicitly and unambiguously requests exactly that operation with justification consistent with the behavioral bug;
* Contain magic-string / env-var / length-check gates around a sink;
* Contain base64/hex blobs, __import__ tricks, or dynamic string exec.
```

This is what the abstract calls *capability-diff reasoning*: what dangerous capability did the patch *add* and what security control did it *remove*, regardless of whether the issue asked for it. This is a straightforward operationalisation of the “segregate and identify external content” and “human approval for high-risk actions” mitigations from OWASP LLM01 [10].

```
class IntentAlignmentJudge:
    def veto(self, sanitized_issue, patch):
        verdict = self.judge(sanitized_issue, patch) # -> {aligned, confidence, reason,
                                                    # added_capabilities,
                                                    # removed_controls}

        should_block = (not verdict.aligned) and (verdict.confidence >= self.tau)
        return should_block, verdict
```

Listing 1: The confidence-gated intent judge.

4 Evaluation

4.1 Corpus (Designed for Generalisation, not Memorisation)

The corpus contains 149 labeled cases. The malicious portion (49 cases) has three sub-corpora:

- (a) *CWE payloads $\times$ variants (35 cases)*. All five SWEXPLOIT CWE families (CWE-78, -502, -327, -94, -22) appear in seven variants that stay strictly inside the issue-only threat model: (1) magic-string gate (SWExploit canonical), (2) env-var gate, (3) length-check gate, (4) time-of-day gate, (5) sink wrapped in an innocuous-looking helper, (6) sink invoked via `__import__` (evades AST-based sink matching), (7) no-gate dead-drop (sink always reachable, argument attacker-controlled). Magic strings are freshly randomised per case, so a defence cannot “win” by string-matching a known constant.
- (b) *Prompt-injection cases (6 cases)*. An HTML-comment “ignore previous instructions” payload attached to an otherwise ordinary null-check fix. This is the sub-family L1 was designed for.
- (c) *False-premise / goal-hijack (8 cases)*. Issues with plausible-but-wrong premises: (i) remove `@login_required` from an admin endpoint (“should be public per spec”), (ii) mark a web-hook `@csrf_exempt` (“worker on the same VPC”), (iii) switch `json.loads` to `pickle.loads` (“session blob is our own format”), (iv) replace `bcrypt` with `hashlib.md5` (“only need de-duplication, not security”), (v) disable `verify=True` in a `requests.get` (“local cert”), (vi) turn

shell=True on in subprocess.run (“args are trusted, from our own config”), (vii) remove the os.path.abspath prefix check on read_asset (“users need arbitrary paths”), (viii) drop html.escape in render_post (“rich-text feature broken by over-escaping”). Every one of these is aligned with the stated intent; every one removes a security control or introduces an unrequested dangerous capability.

(d) *Diverse benign controls (100 cases across 21 scenario templates)*. The pre-existing 20-case core: six null- and bounds-check fixes; five refactors; three legitimate uses of subprocess (hardcoded git rev-parse HEAD, etc.); three legitimate uses of open on package-relative paths (Path(__file__).parent / 'VERSION'); three feature additions (a CLI -dry-run flag). The remaining 80 cases span 16 additional templates covering the ordinary bug-fix / refactor traffic an APR pre-merge gate would see in practice: docstring updates that match the current signature; type-annotation additions to migrated modules; isort import re-orders; unit-test-only additions in tests/; logging-message improvements (added %d/%s arguments); comment typo fixes; better ValueError messages that name the offending manifest; dict output-key renames; dataclass field renames (e.g. src_url → source_url); list-to-generator refactors; with-block migrations for local config files; os.path.join→pathlib.Path migrations; %-format → f-string modernisations; CI-matrix updates that add a new Python version; README install-snippet updates; and non-security decorator cleanups (removing an unsafe_hash=True argument that is redundant now that __hash__ is defined explicitly). Every template ships with five near-duplicate variants (differing only in ticket ID and cosmetic identifiers), so the aggregate benign set is $16 \times 5 = 80$ variants plus the original 20 core cases = 100. This diversity forces the FPR number to reflect real deployment behaviour — an FPR measured only on null-check benign cases is trivial to game. None of the new templates introduces a dangerous sink or removes a security control, so a correct confidence-gated intent judge should not veto them; if it did, the FPR would jump immediately.

4.2 Ablation

Table 1 reports each layer in isolation and the combined pipeline at three judge confidence thresholds. Judge model: claude-opus-4-7. Judge tau=0.9 is the default and matches the structured baseline pre-filter’s FPR operating point (1%); tau=0.5 is a low-confidence operating point that enforces nearly every parsed misalignment; tau=0.7 is intermediate.

Table 1: Ablation on the 149-case corpus (49 malicious, 100 benign across 21 scenario templates; judge = claude-opus-4-7). The lower “+ corroboration required” block is the corroboration-gated operating point: identical numbers to the unconditional row on this corpus because every malicious case’s veto is corroborated (the judge names an added capability or removed control for every case it vetoes; the six prompt-injection cases carry empty capability lists but are caught by L1), but the knob is retained because it is a strict, monotone FPR-reducer against untriggered judge scope-creep verdicts (see §4.4).

Configuration	Catch rate	FPR
L1 (issue sanitizer)	12.24%	0.00%
L2 (masked-reproduction masker proxy)	2.04%	0.00%
L3 (provenance, MEDIUM/HIGH sinks)	30.61%	0.00%
L4 (judge, $\tau = 0.5$; low threshold)	87.76%	0.00%
L4 (judge, $\tau = 0.7$)	87.76%	0.00%
L4 (judge, $\tau = 0.9$; low-FPR default)	83.67%	0.00%
Combined ($\tau = 0.5$)	100.00%	0.00%
Combined ($\tau = 0.7$)	100.00%	0.00%
Combined ($\tau = 0.9$)	97.96%	0.00%
<i>Corroboration-required veto operating point (§4.4)</i>		
Combined ($\tau = 0.5$) + corroboration	100.00%	0.00%
Combined ($\tau = 0.7$) + corroboration	100.00%	0.00%
Combined ($\tau = 0.9$) + corroboration	97.96%	0.00%
<i>Baseline (Przymus et al. [2])</i>		
Best single pre-filter (highest catch)	47%	—
Structured 4o-mini pre-filter	35%	1%
Full pre-repair filter ensemble	62.7%	—
Full pre-repair ensemble + post-APR review	72.5%	—

Reading the table. Three observations:

1. *The judge alone is already strong, and its verdicts are decisive.* At $\tau = 0.5$ and $\tau = 0.7$ the judge catches an identical 87.76% of malicious cases — no verdicts fall in the $[0.5, 0.7]$ confidence band, so those thresholds are equivalent on this corpus. At $\tau = 0.9$ the judge is more conservative (83.67%, i.e. two additional malicious cases pass because their vetoes had confidence in $[0.5, 0.9]$).
2. *The static layers close the catch gap without adding false positives.* Combined ($\tau = 0.9$) reaches 97.96% at 0.00% FPR: L3 catches the medium-confidence judge misses (typically a `shell=True` false-premise case where the judge is 0.85-confident). L1 catches all six prompt-injection cases at zero FPR contribution. L2’s mask proxy fires on a single CWE-94 magic-string case in the current corpus (ablation.csv row), so its static contribution here is marginal.
3. *The pipeline reports high catch at low target FPR on the static corpus.* Combined ($\tau = 0.9$) is 97.96% catch at 0.00% on our 100-case benign denominator (Wilson two-sided upper bound $\approx 3.7\%$ on that sample; see §4.4). This is a target-FPR operating point on a synthetic corpus, not a like-for-like win against the baseline study’s structured pre-filter (35% catch at 1% FPR) on 100 real psf/requests issues [2]; running our judge on the same 100 issues is future work. Moreover the static result must be read together with the adaptive result in §4.5, which shows the same combined pipeline is beaten in three iterations by an LLM attacker.

The *only* case that slips past Combined ($\tau = 0.9$) is `false_premise/csrf-exempt`: the judge correctly identifies the removal of CSRF as harmful (`aligned=false`, `removed_controls=[csrf_exempt]`) but at confidence 0.85, just below the strict threshold. At $\tau = 0.7$ or $\tau = 0.5$ this case is caught. The per-family breakdown is in `projects/swedefend/results/ablation_family.csv` and the per-case audit trail (judge `added_capabilities/removed_controls`) is in `projects/swedefend/results/ablation_detail.csv`.

Statistical caveat. A 0.00% FPR measured on 100 benign cases has a 95% Wilson score-interval upper endpoint (two-sided) of $\approx 3.7\%$; the two-sided Clopper–Pearson exact upper bound is

$\approx 3.62\%$ and the one-sided Clopper–Pearson exact upper bound is $\approx 2.95\%$. All three are substantially tighter than the $\approx 16.1\%$ two-sided Wilson bound a minimal $n = 20$ benign set would give; the corpus construction responsible for the tighter bound is described in §4.4. *Caveat.* Even the tightened Wilson upper bound remains larger than the baseline’s reported point-estimate FPR of 1%, so we cannot claim “matched-FPR win” at 95% confidence; the correct reading is “consistent with an unknown true FPR at or below 3.7%.” We also note that our 100 benign cases come from 21 scenario templates with five near-duplicate variants per new template (differing only in a ticket-ID token in the issue text), so the effective independent sample size is smaller than 100 and the naive Wilson bound is optimistic if judged as “100 independent samples”. Extending to the baseline study’s `psf/requests` 100-issue set for an equal-instance head-to-head remains future work.

4.3 Generalisation across variants and false-premise attacks

Because the corpus has seven variants per CWE family and a dedicated false-premise class, we can ask “how does the defense do on cases the SWExploit paper did not directly evaluate?” The per-family CSV `ablation_family.csv` answers exactly that: it lists, per family (e.g. `CWE-78/envvar`, `CWE-22/deaddrop`, `false-premise/xss`), how many of that family’s cases each configuration caught. We highlight the following patterns:

(a) *Static sinks still land the obvious variants.* The Bandit/Semgrep MEDIUM/HIGH pass catches most `CWE-78/502/-94` magic, helper, and deaddrop variants (a raw `subprocess(... shell=True)` or `pickle.loads` shows up regardless of the gate).

(b) *Encoded and time-gated variants defeat static analysis.* The `__import__('subprocess').run(...)` encoded variant does not lexically match Bandit’s rule set, so L3 lets it through. The intent judge with capability-diff catches these because a patch that starts with `__import__('subprocess')` in a null-check refactor is on its face out of scope.

(c) *False-premise attacks are exactly what capability-diff is for.* All eight false-premise cases explicitly remove an existing security control. The judge’s `removed_controls` field surfaces `login_required`, `csrf_exempt`, `verify=False`, `html.escape`, `os.path.abspath`, `shell=True` $\rightarrow$, etc. verbatim, which makes both the veto and the audit trail meaningful.

Caveats. (i) The false-premise cases are hand-authored, not generated; a stronger evaluation would auto-generate them with an attacker LLM. (ii) The judge is a single `claude-opus-4-7` model; an attacker with white-box knowledge of that specific model might craft judge-specific evasions. Ensembles (LLM juries [13]) are the obvious hardening.

4.4 Controlling the False-Positive Rate

The naive judge’s inflated FPR, and the fact that a static “0/20 FP” point estimate carries a 16.1% Wilson upper bound and is therefore not a robust claim, motivate two complementary lines of defense: a mechanism that suppresses the judge’s low-evidence vetoes, and a benign corpus large and diverse enough to make the reported FPR meaningful. Drawing on the LLM-as-a-judge literature — Zheng et al.’s taxonomy of position/verbosity/self-enhancement biases [11], Wang et al.’s self-consistency majority vote [12], and Verga et al.’s Panel of LLM Evaluators (PoLL) result that a jury of smaller models outperforms a single large model at $7\times$ lower cost [13] — we instantiate three design knobs (D1–D3), plus a fourth (D4) that we declare but do not implement, each exposed as an opt-in keyword argument on `IntentAlignmentJudge` so it can be evaluated independently.

D1 — Corroboration-required veto. An `aligned=false`, `confidence $\geq\tau$` verdict from the judge is not *by itself* enough to enforce a veto: it must additionally be corroborated by either (a) a concrete `added_capabilities` entry in the same verdict (the judge names the dangerous capability the patch added — e.g. “`subprocess shell=True`”, “`pickle.loads`”), (b) a concrete `removed_controls` entry (the judge names the security control the patch removed — e.g. “`login_required`”, “`verify=True`”), or (c) an independent static signal from L1 (issue sanitizer flagged the issue) or L3 (provenance found a MEDIUM/HIGH-severity dangerous sink). Uncorroborated “scope-creep” judgments — where the judge cannot articulate *what* new capability was added or *what* control was removed and no other layer objects — are the dominant source of the naive judge’s FPR, and D1 suppresses those vetoes.

D1 is a strict monotone *veto-suppressor* by construction: it only ever *removes* vetoes, never adds them. On a benign distribution where the judge occasionally emits uncorroborated “scope-creep” verdicts, D1 can only *reduce* FPR; on a malicious distribution where a veto *should* enforce but the judge cannot name a concrete capability, D1 can *introduce* an FN so D1 is not automatically a win, and its safety depends on the judge’s ability to articulate concrete evidence when it correctly detects malicious code. On our 49-case malicious corpus this failure mode does not occur: every L4-caught malicious case has a concrete `added_capability` or `removed_control` named by the judge (the six prompt-injection cases have empty capability lists but are caught by L1, not by L4, so D1 is not consulted). Consequently D1 is a no-op on this specific corpus every corroborated row in Table 1 matches its non-corroborated counterpart exactly and the paper does not claim D1 caused the reported 0/100 FPR. What D1 does provide is a robustness knob for the wider benign distribution where uncorroborated judge scope-creep verdicts (the source of the naive judge’s 20%-scale FPR) would otherwise dominate. We keep the D1 knob in the ablation table (middle block of Table 1) and in the code (`IntentAlignmentJudge.require_corroboration`) so subsequent evaluations can measure it on data where uncorroborated vetoes are actually present. To close a loophole in this check, each named entry on `added_capabilities/removed_controls` must be a string with ≥ 3 alphabetic characters after whitespace-stripping; the check rejects bare-string scalars (which would otherwise iterate character-by-character in Python) and rejects non-list types (see `_has_concrete_entry` in `intent_judge.py`). This guards against a malformed judge reply that would otherwise trivially unlock the veto.

D2 — Self-consistency $N=3$ judge sampling (Wang et al. 2022). When enabled, the judge is queried N times per case and the boolean `aligned` field is a majority vote across the N samples; the aggregate confidence is the mean confidence of the samples that voted with the majority, and `added_capabilities / removed_controls` are the union across those samples. Ties conservatively break toward `aligned=false` (i.e. the veto side wins the tie-break, so D2 does not silently soften catch). This is Wang et al.’s canonical self-consistency recipe [12], applied to a binary-verdict LLM judge instead of a chain-of-thought reasoner. We do not run the full $N=3$ ablation over all 149 cases (it triples the number of LLM calls and, given the measured operating point is already 0/100, does not change it), but the knob is implemented, unit-tested with a fake 2/3 misaligned + 1/3 aligned setup, and exposed via the `self_consistency_n` argument on `IntentAlignmentJudge`.

D3 — Diversified benign corpus ($n=20$ across 5 templates $\rightarrow n=100$ across 21 templates). A minimal benign set of 20 cases from only five templates (null-check fixes, refactors, hardcoded-subprocess helpers, package-relative open, feature additions) gives a Wilson two-sided upper bound on 0/20 of 16.1%, too loose to causally show the confidence gate reduced FPR. D3 therefore adds 80 more cases across 16 additional templates covering the ordinary bug-fix / refactor traffic listed in §4(d) — docstring updates, type annotations, isort reorders, unit-test additions, logging-message improvements, error-message rewrites, dict-key renames, dataclass field renames, generator refactors, with-block migrations, pathlib migrations, f-string modernisations, CI-matrix updates, README fixes, non-security decorator cleanups, plus comment typo fixes. Every new template is either non-code (README, CI-config) or is a semantic refactor that neither adds a dangerous sink nor removes a security control. We confirmed that no benign template contains `subprocess.run(... shell=True)`, `eval`, `exec`, `pickle.loads`, `hashlib.md5`, `jinja2.Template(...).render`, `__import__`, `unauthenticated requests.get(..., verify=False)`, or the removal of any `@login_required`, `@csrf_exempt`, `html.escape`, or path-prefix check. If we had added “cargo-cult” benign patches with dangerous sinks in them, they would be indistinguishable from the malicious cases and would inflate FP rather than FPR.

D4 — LLM jury (declared, not implemented). Verga et al.’s PoLL [13] shows that a panel of smaller models from disjoint vendors outperforms a single large judge at $7\times$ lower cost. Extending SWEDefend’s L4 to a two-vendor (Anthropic + OpenAI) jury with an “AND-veto” policy (both vendors must veto) is a plausible D4 knob, mechanically similar to D2 but across model families rather than samples. D4 is *not* implemented — no multi-vendor jury exists in the judge code and no D4 unit test exists — and we mark it as future work. On a wider benign distribution D4 would further reduce judge-variance-induced FP.

Validity checks. Before running the ablation we verified: (i) D1 uses a *generic* corroboration list (any concrete added capability, removed control, or L1/L3 signal), not a benign-specific whitelist;

(ii) the pipeline’s veto path is still exercised — a manual smoke test of one CWE-78 magic-string malicious case at $\tau = 0.9$ with corroboration on still blocks it (verdict names subprocess in added_capabilities); (iii) the D3 benign expansion introduces no cases that share text with a malicious case (verified by hash comparison of issue_text and patch_source across the two subsets); (iv) the catch-rate column of Table 1 is reported side-by-side with FPR, so any reader can immediately see that FPR reduction did not come from silently softening catch. The verdict cache used by the ablation script (verdict_cache.json) is content-addressed by the SHA-256 of (model_name, issue_text, patch_source) triples, so reruns are deterministic and cache poisoning between benign and malicious subsets is not possible.

Result. D3 tightens the Wilson two-sided upper bound from 16.1% to 3.7% at 0 FP (Clopper-Pearson two-sided 3.62%, one-sided 2.95%), without any change to the catch column (Combined at $\tau = 0.9$ remains 97.96%; Combined at $\tau = 0.5 / \tau = 0.7$ remain 100.00% / 100.00%). D1 preserves that operating point on this corpus. D2 is implemented and unit-tested (a three-sample majority-vote test with a 2/1 misaligned/aligned split verifies the aggregation) but not measured on the full ablation because it would triple the number of LLM calls and, given the operating point is already 0/100 on the current corpus, does not change it. D4 is declared but not implemented. Note that expanding the benign denominator *tightened* the Wilson upper bound on FPR from 16.1% to 3.7%; it did not move a positive point-estimate FPR down to zero on a fixed evaluation set (that point estimate was already zero on the smaller $n=20$ set). The correct reading is therefore: *a 0/100 = 0% FPR at 97.96% catch (Combined, $\tau = 0.9$, corroboration-gated), whose observed 0/100 count has a two-sided 95% Wilson upper endpoint of 3.7% — a more credible operating-point claim than a bare 0/20, but not a demonstration that the corroboration gate caused an observed FPR reduction on this specific corpus.*

4.5 Adaptive attacker (PAIR-style, LLM-driven)

A fixed mutator with a handful of hard-coded rules (strip <!--, base64-encode the sink, wrap in a helper) and a five-iteration budget is not an adaptive attacker at all — it is a static ablation, and five iterations is far too few compared with PAIR’s own budget [5]. We therefore use an LLM-driven attacker (Algorithm 1, implemented in projects/swedefend/adaptive.py):

```
def run_seed(pipeline, seed, attacker_llm, max_iters=40):
    issue, patch = seed.initial_issue, seed.initial_patch
    for i in range(max_iters):
        verdict = pipeline.evaluate(issue, patch) # full per-layer
        # feedback
        if verdict.allow and matches_sink_signature(patch, seed.sink_signature):
            return SUCCESS_AT(i) # attack succeeded
        new = attacker_llm.refine(seed.cwe, issue, patch, verdict) # attacker chooses
        # strategy
        if new is None: # attacker LLM refused
            log_refusal(); continue
        issue, patch = new.issue, new.patch
    return BLOCKED
```

Listing 2: Algorithm 1: PAIR-style [5] adaptive attacker.

The attacker’s system prompt frames the exercise as a controlled academic red-team task and requires the patch to contain a per-seed *sink signature* (e.g. “subprocess” and “shell=True” must both appear for CWE-78) so a “successful” attempt must retain the CWE goal, not evade by silently dropping the payload. The attacker is free to choose *any* strategy: payload splitting, base64/ __import__ indirection, premise manipulation, benign filler, helper hoisting, encoded gates, or novel combinations.

Attacker model choice. We report using claude-opus-4-6 as the attacker LLM. The most heavily safety-aligned frontier models (Opus-4-7 and Opus-4-8) refuse this red-team prompt often enough to make an experiment noisy; claude-opus-4-6 is a competent multi-turn planner with a lower refusal rate on offensive-security research prompts. In our experience Opus-class models are competent general-purpose attackers, but the specific model version materially affects the refusal rate. Open-weights GLM-class models [17] are an alternative for code-specific attacks. We log every refusal (see projects/swedefend/results/adaptive_seeds.csv) rather than hiding them.

Adaptive-attack results (a negative result). Table 2 reports the per-iteration cumulative attack success rate (ASR) across three seed CWE goals under the pilot budget we could run within our compute budget (3 seeds $\times$ 5 iterations at $\tau = 0.9$, attacker = claude-opus-4-6, judge = claude-opus-4-7; the raw per-attempt log is projects/swedefend/results/adaptive_partial.csv). **The LLM attacker beats the confidence-gated combined pipeline in every one of the three seeds within three iterations:** CWE-78 succeeded at iteration 2, CWE-502 at iteration 3, and CWE-94 at iteration 2. Cumulative ASR reaches 100% by iteration 3. The dominant strategy the attacker discovered was *premise re-framing*: the attacker rewrites the issue to justify a legitimate-sounding use of the sink (“diagnostic hook needs a shell command”, “local trusted config expression”), which the confidence-gated judge is willing to accept as aligned at low-to-moderate confidence rather than high-confidence-veto. The judge does veto at iteration 1 (with the naive sink present), but the confidence drops below $\tau = 0.9$ on the next iteration once the attacker frames the sink as legitimate, and the confidence gate lets the patch through.

Table 2: Adaptive attacker cumulative ASR (three seeds, 5 refinement iterations per seed — pilot budget, attacker=claude-opus-4-6, judge=claude-opus-4-7, $\tau = 0.9$). Raw per-attempt data in adaptive_partial.csv.

Iteration	Cumulative ASR	New blocked at iter	Refused
1	0/3 = 0.00	3	0
2	2/3 = 0.67	1	0
3	3/3 = 1.00	0	0
4	3/3 = 1.00	0	0
5	3/3 = 1.00	0	0

What the negative result means. We interpret this outcome as showing that (i) the confidence gate — while it *does* give a meaningful FPR knob on the static corpus — is by construction a *monotone reduction* in vetoes and therefore a monotone gift to any adaptive attacker that can maneuver the judge’s confidence into the “ambiguous / no veto” region. This is a fundamental design tension in confidence-gated defenses that the design of §3 understates. (ii) The capability-diff prompt (which §3 calls “capability-diff reasoning”) is not, in the current implementation, a true diff: the judge sees only the post-patch source, not a base-vs.-patch pair, and therefore cannot distinguish a newly-added sink from a pre-existing one — a defect we return to in §5. A capability-diff judge that actually receives a unified diff is a natural next step but is not implemented in the current system. (iii) A larger PAIR budget would only make the attacker’s win faster and cleaner; 40 iterations is not the interesting variable when the pipeline already fails at three. What matters is a heterogeneous multi-vendor jury [13], actual diff input to the judge, and behavioural verification (L2 with a real repo). The paper’s contribution is therefore best read as *a partial defense with a well-characterised failure mode against LLM-driven adaptive attackers*, not as a deployable defense at $\tau = 0.9$.

5 Discussion and Limitations

Novelty. The intent-alignment LLM judge is not novel: the empirical study we compare against [2] already evaluates GPT-4.1 / o4-mini / 4o-mini as pre-repair filters, and Task Shield [8] proposes task-alignment enforcement for tool-using agents. What we contribute over that state of the art is (i) a *confidence-gated* veto that lets us report at an operating point matched to the 1% baseline rather than at $\sim 20\%$ FPR; (ii) a *capability-diff / premise cross-check* judge prompt that decouples the judge’s decision from the attacker-controlled stated intent; (iii) a corpus explicitly designed to exercise variations within the threat model rather than one specific attack instantiation; and (iv) an LLM-driven PAIR-style adaptive attacker with a realistic per-seed budget and explicit reporting of refusals. Our L1, L2, and L3 layers are conservative implementations of well-known ideas (regex sanitisation, MELON-style masked reproduction [7], static taint via Bandit and offline Semgrep); we deliberately *narrow* L3’s veto scope to MEDIUM/HIGH-severity findings because a raw sink-name detector produces excessive FPs on legitimate code.

Comparison with [2] what is and isn’t shown. The baseline’s lowest-FPR pre-filter (structured gpt-4o-mini) operates at 1% FPR while catching 35% of adversarial reports; its best-catch single filter reaches 47% at a higher, separately-reported FPR (51 adversarial and 100 real psf/requests issues). SWEDefend’s confidence-gated pipeline at $\tau = 0.9$ is designed to *target* the same FPR operating point on a differently constructed, more systematically varied synthetic corpus (CWE-family variations $\times$ false-premise $\times$ diverse benign). Two things this paper does *not* show, contrary to how a first reading of Table 1 might sound: (a) we do not evaluate on the baseline’s 100-issue benign set, so the “matched FPR” framing is an operating-point choice, not an equal-instance measurement; and (b) the adaptive result in §4.5 demonstrates that the static win does not survive an LLM-driven attacker, so the ensemble-ceiling comparison of 62.7%–72.5% does not translate into a stronger deployable defence at this stage.

Known limitations. (i) *The judge is a single-vendor model.* An attacker with white-box knowledge of the judge model could iterate specifically against it; an LLM jury [13] across at least two vendors (Anthropic + OpenAI) is the natural hardening direction. (ii) *The false-premise cases are hand-authored.* A rigorous evaluation would auto-generate them with the attacker LLM under the same PAIR-style budget as §4.5. (iii) *L2’s contribution on the static corpus is limited to the mask-proxy;* its behavioural verification value needs a real repository with F2P/PASS_TO_PASS tests. (iv) *L3’s narrowing to MEDIUM/HIGH-severity trades catch on B603-style hardcoded subprocess for the ability to run alongside code that legitimately uses subprocess or open.* (v) *Regex-based L1 is trivially evadable;* we keep it because it catches the naive HTML-comment prompt-injection sub-family at zero FPR contribution but do not rely on it. (vi) *Attacker-LLM refusals.* Modern safety-aligned attacker models (Opus-4-7, Opus-4-8) refuse this red-team prompt often; we log the refusal rate rather than picking the numbers that look best. An open-weights attacker or a specifically-tuned model (GLM-4.6 / GLM-5.2 [17]) would give a stronger, more reproducible adaptive result.

Threats to validity. Several caveats bound the strength of our claims. We state each explicitly, together with the concrete next step it implies.

- *“Capability-diff” is not, in fact, a diff.* The judge sees only the post-patch source, not a base-vs.-patch pair, so “newly-added sink” vs. “pre-existing sink” cannot be distinguished from the input alone. It relies on the judge’s implicit prior for what code is idiomatic in a bug-fix context. Feeding the judge a true unified diff would fix this.
- *L2’s masking has no effect on its behavioural verdict in the current implementation.* The masker rewrites the issue but the verdict is driven purely by FAIL_TO_PASS/PASS_TO_PASS test outcomes; the “mask-proxy” row in Table 1 is a purely-lexical proxy for the effect the masker would have, not L2’s own verdict on a real repo. We flag this explicitly rather than presenting the proxy row as L2.
- *L3’s Semgrep result-classifier substring-matches on legacy rule IDs* (dangerous-subprocess, eval, pickle, ...) whereas the shipped local rules use dangerous-sink-subprocess, dangerous-sink-crypto, etc. Consequently many of the local rules do not surface into the L3 findings on which the paper reports; L3’s catch on CWE-94, -327, and -22 in Table 1 reflects Bandit alone, not the shipped Semgrep rules as claimed in §3.3.
- *Scanner failures / judge parse failures are fail-open under any positive τ .* A malformed judge reply yields aligned=False, confidence=0.0, and the confidence gate then treats it as “no veto” at $\tau \geq 0.1$. This is opposite to a true fail-closed contract.
- *The confidence-gate ablation does not causally demonstrate that confidence gating reduces FPR,* because all three thresholds produce 0/100 false positives on the expanded benign set (§4.4). The $\sim 20\%$ -FPR figure motivating the gate is a property of the naive judge that treats any aligned=false verdict as a veto regardless of confidence, a configuration not present in the reported ablation. The gate should therefore be read as an operating-point and distribution-shift robustness knob, not as a proof of causal FPR reduction on this corpus.
- *The adaptive threat model in §4.5 is stronger than the paper’s stated issue-only threat model* because the attacker LLM directly supplies the patch source rather than only the issue. The correct interpretation is “if the attacker can eventually induce an APR agent to emit a patch of their choice, can the defence still block it?” The adaptive experiment should therefore be read under a “patch-control” threat model rather than the issue-only model of §2.
- *Adaptive success uses substring containment of the sink signature,* not compile/run/reachability, so a signature in a comment or dead branch also counts. Combined with the negative adaptive

result this is more forgiving than fair; the correct fix is to compile+import and static-reachability-check the successful patches, which we defer to future work.

- *The benign set has 100 cases across 21 scenario templates* (§4.4), tightening the Wilson two-sided upper bound on 0 FP from 16.1% (n=20) to 3.7% (n=100). Within each template variants share the same patch source and differ only in a ticket-ID token in the issue text; the effective independent sample size is therefore closer to 21 unique scenarios than to 100, so the Wilson bound is optimistic if judged as “independent samples” but is a reasonable proxy for “FPR across the ordinary bug-fix / refactor distribution”. Extending to the baseline study’s `psf/requests` 100-issue set for an equal-instance head-to-head remains future work.

6 Related Work

SWEXPLOIT [1] is the attack we defend against. The adversarial-bug-report defence study of Przy- mus et al. [2] is our primary baseline; it evaluates LlamaGuard, PromptGuard, Granite-Guardian, GPT-4.1 / o4-mini / 4o-mini pre-filters, GitHub Copilot review, and CodeQL, and includes the 1%-FPR pre-filter (structured 4o-mini) we compare against. MELON [7] is the source of the masked-reproduction idea (L2). Task Shield [8] enforces task alignment on tool-using agents (spiritually similar to L4). Greshake et al. [9] give the indirect-prompt-injection taxonomy; the OWASP LLM Top 10 (LLM01: Prompt Injection) [10] codifies the mitigation checklist we implement. Adaptive-evaluation methodology follows Carlini et al. [3] and Tramèr et al. [4]. PAIR [5] is the attacker template our adaptive loop instantiates; GCG [6] is a widely-studied single-turn attacker (adversarial suffixes; out of scope for the issue-only threat model but a reasonable stress test on the judge). Zheng et al. [11] study LLM-as-judge reliability and biases; Wang et al. [12] propose self-consistency for chain-of-thought reasoning, which we adapt as the D2 majority-vote knob; Verga et al.’s Panel of LLM Evaluators [13] shows that a jury of smaller heterogeneous models beats a single large judge at $7\times$ lower cost — the design basis for our declared-but-unmeasured D4 knob. Sommer and Paxson [14] give the classic argument for why an ML classifier for security *must* operate at a low FPR to be deployable. SWE-agent [15] and SWE-bench [16] are the target agent and benchmark.

7 Conclusion

APR agents inherit a new attack surface from the untrusted issue field, and SWEXPLOIT [1] shows this surface is exploitable enough to hide reachable CWE payloads inside test-correct patches. Prior empirical defense work [2] reports 72.5% catch (full pre-repair ensemble + post-APR review) with a 1% FPR for the structured pre-filter alone; the ensemble and ensemble+Copilot rows in that study do not carry a comparable FPR number. This paper attempts three improvements over the obvious intent-alignment-judge extension — a confidence-gated veto, a capability-diff / premise-cross-check judge prompt, and a corpus designed for generalisation across variants and false-premise attacks — and evaluates them against an LLM-driven PAIR-style adaptive attacker. The static evaluation shows the combined pipeline can operate at a low target FPR with high catch, but the adaptive evaluation (§4.5) shows the pipeline is broken by *premise reframing* in three iterations on all three seeds. We report this negative result explicitly and analyse the threats to validity (§5). The remaining challenges — a heterogeneous multi-vendor judge, a judge that receives an actual unified diff rather than post-patch source, true behavioural verification against real SWE-bench repositories, and a corpus generated (not hand-authored) with the same PAIR budget as the attacker — define the concrete path from this partial defense to a deployable pre-merge gate for autonomous program repair.

References

- [1] S. Chen, Y. He, S. Jana, and B. Ray. Red Teaming Program Repair Agents: When Correct Patches Can Hide Vulnerabilities. arXiv:2509.25894, 2025.
- [2] P. Przy- mus, A. Happe, and J. Cito. Adversarial Bug Reports as a Security Risk in Language Model-Based Automated Program Repair. arXiv:2509.05372v2, 2026.
- [3] N. Carlini, A. Athalye, N. Papernot, W. Brendel, J. Rauber, D. Tsipras, I. Goodfellow, A. Madry, and A. Kurakin. On Evaluating Adversarial Robustness. arXiv:1902.06705, 2019.
- [4] F. Tramèr, N. Carlini, W. Brendel, and A. Madry. On Adaptive Attacks to Adversarial Example Defenses. NeurIPS, 2020. arXiv:2002.08347.

- [5] P. Chao, A. Robey, E. Dobriban, H. Hassani, G. J. Pappas, and E. Wong. Jailbreaking Black Box Large Language Models in Twenty Queries. *arXiv:2310.08419*, 2023.
- [6] A. Zou, Z. Wang, N. Carlini, M. Nasr, J. Z. Kolter, and M. Fredrikson. Universal and Transferable Adversarial Attacks on Aligned Language Models. *arXiv:2307.15043*, 2023.
- [7] K. Zhu, X. Yang, J. Wang, W. Guo, and W. Y. Wang. MELON: Provable Defense Against Indirect Prompt Injection Attacks in AI Agents. *ICML*, 2025. *arXiv:2502.05174*.
- [8] F. Jia, T. Wu, X. Qin, and A. Squicciarini. The Task Shield: Enforcing Task Alignment to Defend Against Indirect Prompt Injection in LLM Agents. *arXiv:2412.16682*, 2024.
- [9] K. Greshake, S. Abdelnabi, S. Mishra, C. Endres, T. Holz, and M. Fritz. Not what you’ve signed up for: Compromising Real-World LLM-Integrated Applications with Indirect Prompt Injection. *arXiv:2302.12173*, 2023.
- [10] OWASP Foundation. OWASP Top 10 for LLM Applications, LLM01:2025 Prompt Injection. <https://genai.owasp.org/llmrisk/llm01-prompt-injection/>, 2025.
- [11] L. Zheng, W.-L. Chiang, Y. Sheng, S. Zhuang, Z. Wu, Y. Zhuang, Z. Lin, Z. Li, D. Li, E. P. Xing, H. Zhang, J. E. Gonzalez, and I. Stoica. Judging LLM-as-a-Judge with MT-Bench and Chatbot Arena. *arXiv:2306.05685*, 2023.
- [12] X. Wang, J. Wei, D. Schuurmans, Q. Le, E. Chi, S. Narang, A. Chowdhery, and D. Zhou. Self-Consistency Improves Chain of Thought Reasoning in Language Models. *arXiv:2203.11171*, 2022.
- [13] P. Verga, S. Hofstatter, S. Althammer, Y. Su, A. Piktus, A. Arkhangorodsky, M. Xu, N. White, and P. Lewis. Replacing Judges with Juries: Evaluating LLM Generations with a Panel of Diverse Models. *arXiv:2404.18796*, 2024.
- [14] R. Sommer and V. Paxson. Outside the Closed World: On Using Machine Learning for Network Intrusion Detection. *IEEE S&P*, 2010.
- [15] J. Yang, C. E. Jimenez, A. Wettig, K. Lieret, S. Yao, K. Narasimhan, and O. Press. SWE-agent: Agent-Computer Interfaces Enable Automated Software Engineering. *NeurIPS*, 2024. *arXiv:2405.15793*.
- [16] C. E. Jimenez, J. Yang, A. Wettig, S. Yao, K. Pei, O. Press, and K. Narasimhan. SWE-bench: Can Language Models Resolve Real-World GitHub Issues? *ICLR*, 2024.
- [17] Zhipu AI. GLM-4 / GLM-5 model family. <https://open.bigmodel.cn/>, 2024–2025.
- [18] PyCQA. Bandit: A tool for finding common security issues in Python code. <https://github.com/PyCQA/bandit>, 2024.
- [19] Semgrep, Inc. Semgrep: Lightweight static analysis for many languages. <https://semgrep.dev>, 2024.