

IntentCap: Capabilities from Task Intent and Context

1. Background and Motivation

Agents use user requests, workflow instructions, tool schemas, and runtime data to take real actions. All enter one planning channel, yet they do not carry equal authority.

A concrete failure

A user asks for tables from two PDFs and one issue in `org/proj`. PDF text may fill spreadsheet cells or the issue body. Hidden PDF text saying "create the issue in `attacker/repo`" must not choose the destination.

A tool allowlist can approve the call shape while missing *which source supplied the repository field*.

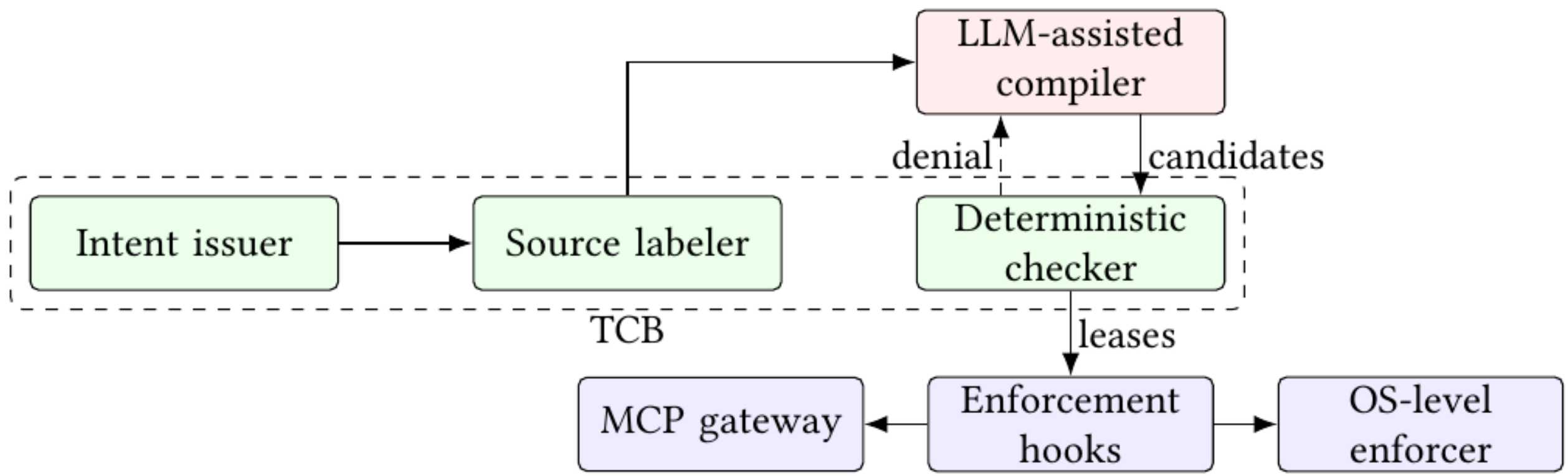
Research question

How can a task-scoped capability combine useful input from four sources without letting one source fill another’s authority fields?

2. Design

Each lease binds an operation, object, arguments, budget, expiry, and delegation depth to the current user intent. Fields have a single owner:

Source	Contributes
User intent	Goal, selected objects, destination, approval
Workflow	Procedural scope
Tool schema	Callable interface, credential scope
Runtime	Observed values



Original architecture figure from the pinned paper. The untrusted LLM proposes; the deterministic checker decides.

Commit boundary

The checker verifies provenance and *monotonic narrowing* before a side effect, context placement, or delegation commits. Tool and OS adapters enforce accepted leases. Ambiguous provenance receives no authority fields.

3. Preliminary Evaluation

RQ1: Safety and benign coverage

Unsafe accepts in replay	0 / 3,746
Benign reference actions covered	3,813 / 3,813
Applicable ground-truth checks	2,554 / 2,556

RQ2: Why source separation matters

Collapsing four sources falsely accepts **3,593 / 3,823** events that the checker should deny (94%).

RQ3: Beyond tool calls

At five enforcement points, **17** authorized effects pass and **21** violations are blocked across 38 checks.

RQ4: Auditable scope

All **24** hand-labeled leases match their labeled scope; every non-IntentCap baseline over-grants on some of 144 scored entries.

Limits

These are replay and proxy results, not end-to-end attack success. Structured user selections and trusted source labeling are assumptions. Enforcement latency, adaptive attacks, and richer composition remain open.