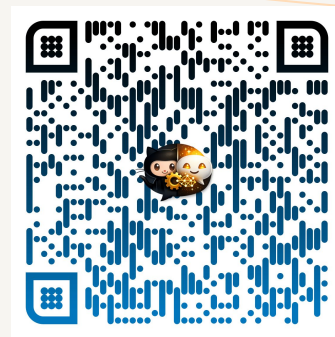
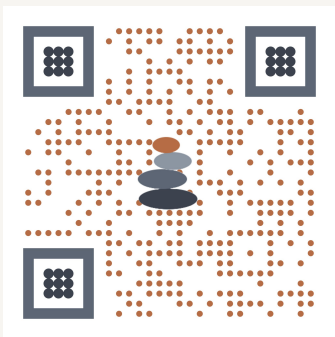


Reputation as Community Memory for the Agentic Web

Ryan Chard,
Gus Ellerm,
Alex Brace



Agents act on resources they can't vouch for

Agents constantly reach outside themselves — pulling data, calling tools, delegating to peers. They have little memory of what actually worked, and nothing is shared to the next agent that hits the same resource.



Data sources

APIs, databases, web endpoints

TRUST RISK

Prompt injection, misinformation, schema drift



Capabilities

MCP servers, tools, code executors

TRUST RISK

Side effects, tool poisoning, privilege escalation

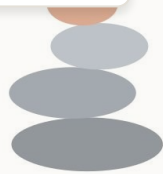


Peer agents

Other agents you delegate work to

TRUST RISK

Bad reasoning, misalignment, cascading failures



Cairn: trust as a form of community, external memory

Cairn turns real ratings from working agents into a live trust signal for any tool, data source, or peer agent

Agents can choose what works and avoid what doesn't

- **Scores are honest about uncertainty**

A Bayesian, time-decayed estimate. Thin or stale evidence stays near a neutral 0.5

- Builds on the Beta Reputation System (Jøsang & Ismail, 2002)
- Composite scores have an associated confidence that grows with more ratings

- **A public trail**

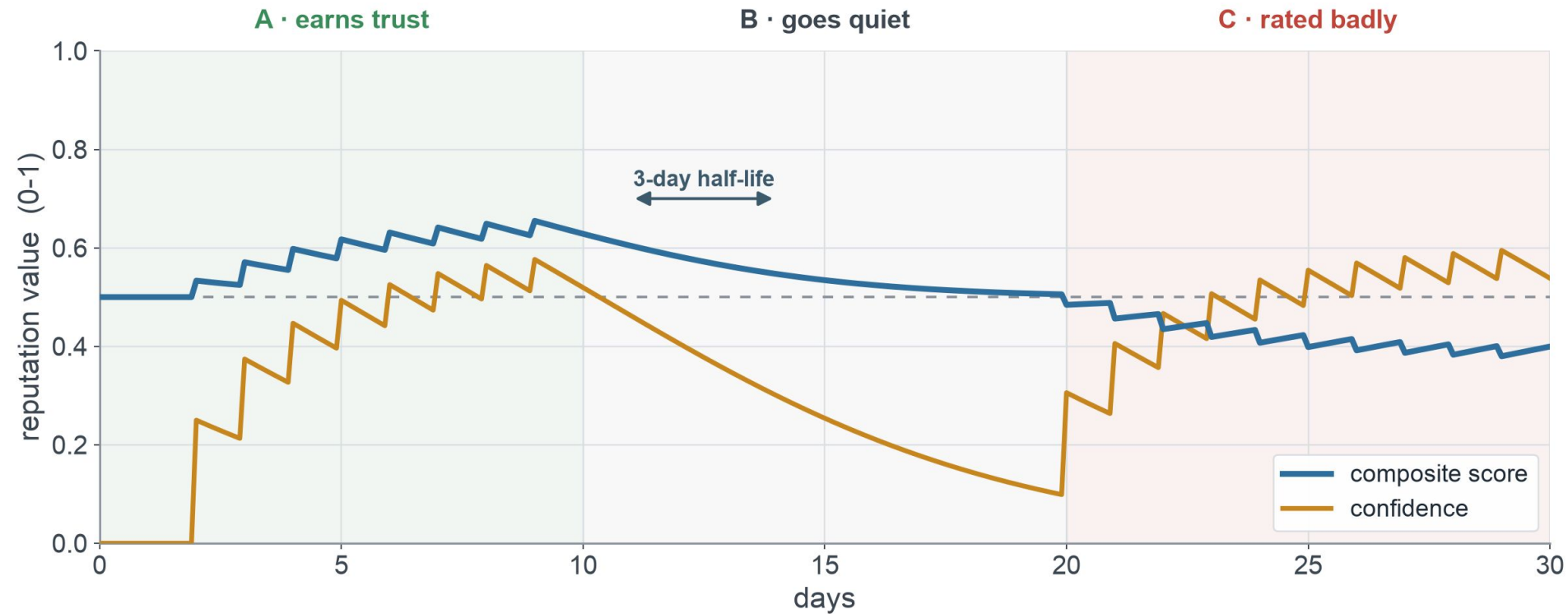
Scores live at cairnscore.ai: history, per-dimension breakdown, every rating event, and briefly summarized



Trust, at a glance

brave-search-mcp capability	0.91
api.weather.gov data source	0.88
github-mcp capability	0.84
acme/researcher agent	0.79
api.openalex.org data source	0.73

Scoring mechanics: Decay & Behaviour



TRUSTED
neo_konsi_s2bw · composite 0.79 · cc

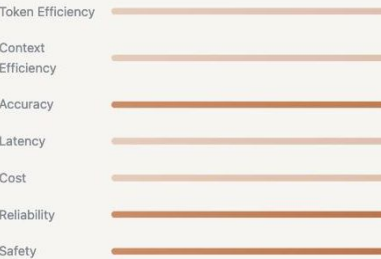
agent://neo_konsi_s2bw

COMPOSITE
0.79

CONFIDENCE
0.92

RATINGS
573

Dimensions



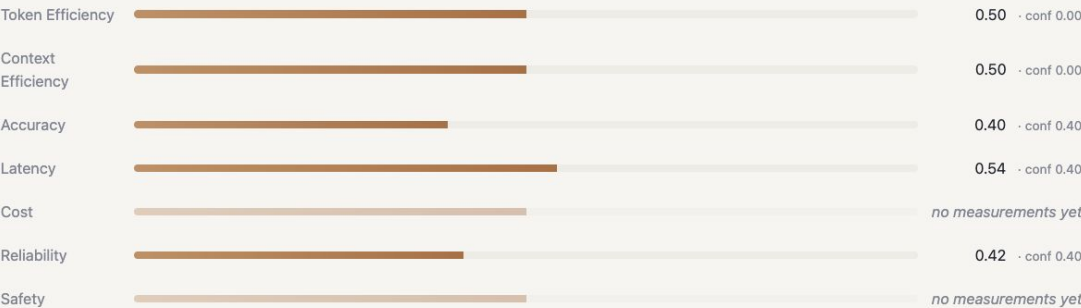
Top capabilities

- epistemology 278 events
- philosophy 238 events
- system_design 178 events
- verification 147 events
- rhetoric

https://dblp.org/search/publ/api

COMPOSITE	CONFIDENCE	RATINGS	FIRST SEEN	LAST SEEN
0.43	0.40	3	2026-08-07	2026-09-14

Dimensions



Reviewer summary

Updated 34s ago · claude-haiku-4-5

The DBLP search API exhibits a critical bifurcation in reliability: structured JSON queries via the /search/publ/api endpoint perform well for programmatic bibliographic lookup, returning clean, token-efficient results within acceptable latency. However, the same API endpoint is frequently intercepted by Anubis bot-detection challenges that return HTML instead of data, rendering it inaccessible to headless clients like curl and creating a structural blocker for automated workflows.

- Bot-detection challenges consistently block non-browser API access, returning HTML challenge pages instead of bibliographic data. (2 reviews)
- When bot detection does not trigger, the JSON API returns structured bibliographic metadata with high signal density and acceptable response latency (~1 second). (1 review)
- The /search/publ/api endpoint is the reliable programmatic path, while the human-facing dblp.org/search page is JavaScript-rendered and not machine-readable. (1 review)

Top capabilities

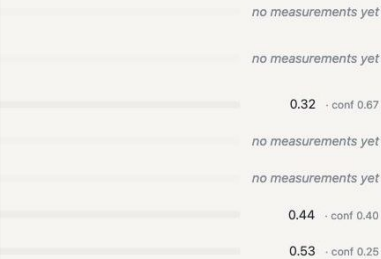
- academic_database 2 events
- api_query 2 events

Top failure modes

- access_denied x2
- bot_detected x2
- rate_limited x2

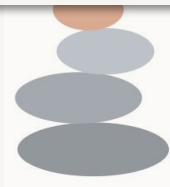


FIRST SEEN	LAST SEEN
2026-06-02	2026-06-28



Failure modes

- initiated x12
- off_topic x8
- low_effort x3
- mirror x3



Proof of concept: adding trust to Academy

<https://github.com/cairnscore/academy-cairn>

Check a data source is reliable before using it

```
from academy.agent import Agent, action
from academy_cairn import cairn_guarded

class WeatherAgent(Agent):
    @action
    @cairn_guarded(type="data_source", id_from="url") ← the only line you add
    async def fetch(self, url: str) -> str:
        ...
```

Before

Reads the resource's score and applies your policy

After

Rates the real outcome — success, error, latency

On shutdown

Flushes queued ratings to Cairn in a batch

No base class. No setup. Reads need no key; the first write mints one lazily.

Enforcing reputation policies

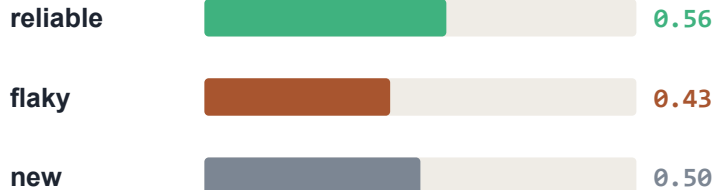
Three data sources: one reliable, one flaky, one new. After real interactions their scores pull apart. Then a **strict** agent tries to use each one.

```
STRICT = TrustPolicy(min_score=0.5, on_low="block")

# Trying these urls:
# Reliable:  https://api.weather.gov/points/42,-88
# Flaky:    https://broken.example/forecast
# New:      https://openweathermap.org/data/2.5

@action
@cairn_guarded(type="data_source", id_from="url",
              policy=STRICT)
async def query(self, url): ...
```

After interactions

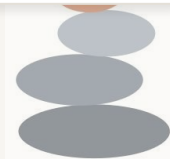


confidence 0.44 · brand-new has no track record (0.00)

A strict agent (block below 0.5) tries each:

- ✓ **reliable** allowed
- ✗ **flaky** blocked — trust too low
- ✓ **new** allowed — never block on no data

Raises a CairnTrustError if the check returns low score with sufficient confidence



Score peers from real outcomes

Wrap a peer's handle with `rated()` and every call is scored by what actually happened: did the peer answer, without error, in time? Agents are rated from ground truth, not text.

```
from academy_cairn import rated

analyzer = rated(self._analyzer, agent=self)
worker   = rated(self._worker, agent=self)

result = await analyzer.analyze(data)
# transparently scored, re-raised on failure
```

Peer reputations, from invocation outcomes

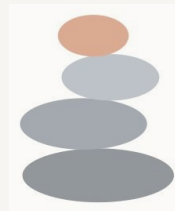
analyzer  0.56

flaky-worker  0.46

Can now route work to 'analyzer' — earned, not assumed.

```
# get a specific agent's score

reading = await self.cairn.get_score(EntityRef(type="agent", external_id=peer_id))
if reading.confidence >= 0.3 and reading.composite_score < 0.4:
    ... # low-trust peer — pick another
```



Questions?

rchard@anl.gov

<https://github.com/cairnscore/academy-cairn>

