

UnitSense



Track the spend. Prove the return.

Every AI dollar — API, subscriptions, coding agents — attributed to a team, project, or person, with limits that actually stop spend.

Bilal Baqar — CEO Salman Sikandar — President & CTO

`bilal@unitsense.ai` • `salman@unitsense.ai`

The tools arrived. The controls didn't.

Multiple providers. API usage and per-seat subscriptions, tracked in different places by different teams. Coding agents on top. The spend is real and growing — the accounting never caught up.

Many providers

API usage here, subscriptions there — no single view of what the company actually spends on AI.

No attribution

The bill says how much. It never says which team, which project, or which person.

Guardrails on paper

If rules exist, they're a policy document asking people to be careful. Nothing enforces them.

Three questions you should be able to answer.

01

Where did it go?

Not the total — the breakdown. Which team, which project, which person, which model. Down to the individual request, if anyone asks.

02

Who decided that?

A spike has an owner. When spend jumps 30% in a month, someone chose the model, shipped the feature, or left the seat idle — and today nobody can say who.

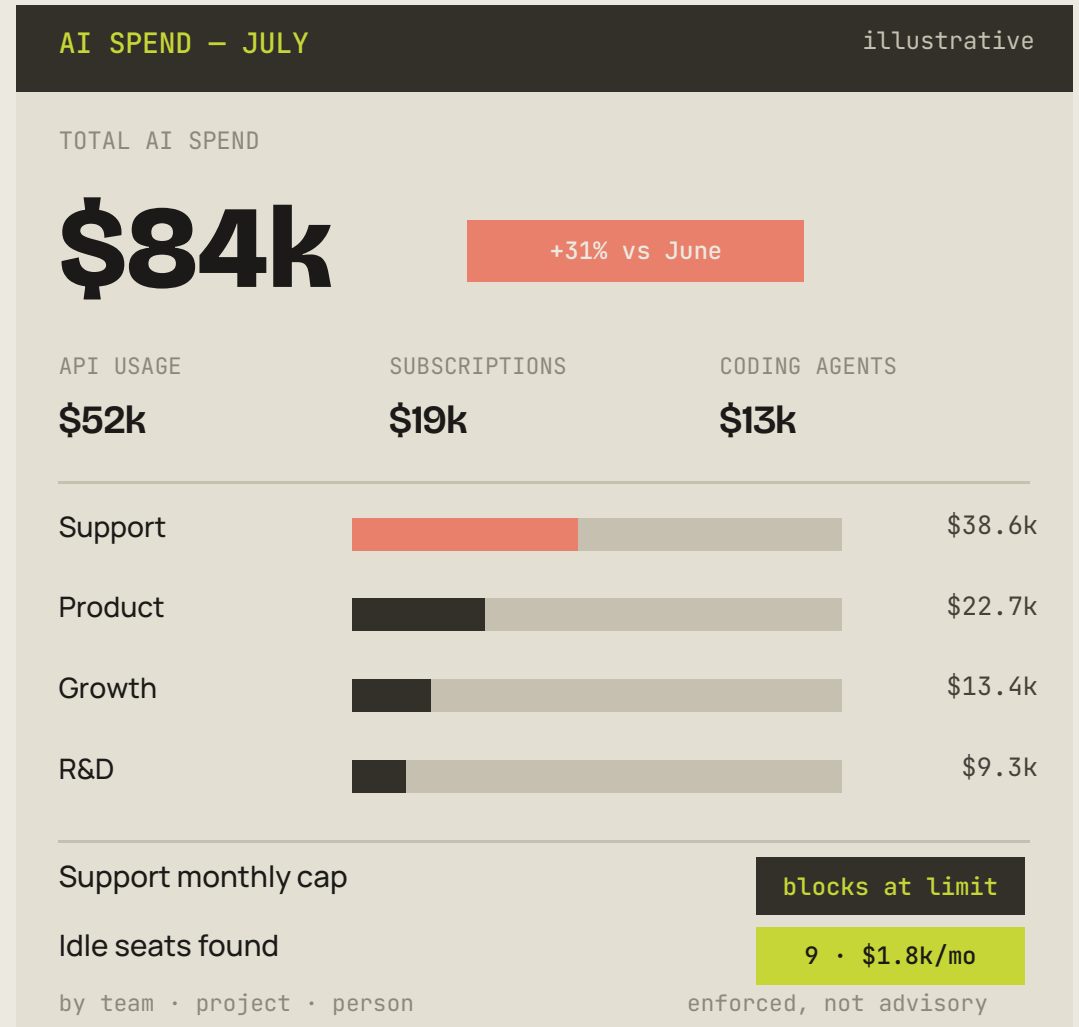
03

Can you stop it?

Alerts tell you it happened. A real answer is a limit that holds — spend that stops at the budget instead of sailing past it into next month's bill.

Every AI dollar, by team, project, and person.

- 01 Spend by team, project, and person — down to the individual request
- 02 API, subscriptions, and coding agents in one view
- 03 Reconciled line-by-line to provider invoices



A budget that can't say no isn't a budget.

Set a limit on any dimension — team, project, model, key, person. Choose alert or block. Block means the request is denied at the limit, in the request path — not an email after the fact.

POLICY — SUPPORT · MONTHLY		illustrative
Limit	\$6,000	
Warn at	80%	
Action	Block	
\$6,000 of \$6,000 used		requests denied

Governance by PDF

The alternative is a policy document asking people to be careful — pages of guidance per role, no meter, no limit. Everyone means well, and the bill arrives anyway. A limit that enforces itself replaces all of it.

ALERT WHEN YOU WANT A NUDGE. BLOCK WHEN YOU MEAN IT.

A per-seat subscription looks like a fixed cost. It isn't.

It's tokens with a flat price on top — so it can be attributed, measured, and right-sized like everything else you buy.

SEATS — CODING TOOLS		illustrative
PROVISIONED	ACTIVE	IDLE
34	25	9
Idle-seat waste		\$1.8k/mo
Value per seat (estimated)		1.4×
Top project by tokens		checkout-agent

The other half of the bill

Most companies' AI spend splits between per-token APIs and per-seat subscriptions — tracked in different places, by different teams. UnitSense reopens the subscription line: who has a seat, who actually uses it, what it's worth (estimated), and where that work goes by project.

Built for the stack you already run.

MODELS	OpenAI · Anthropic · Google Gemini · AWS Bedrock · Azure OpenAI · Mistral
--------	---

VOICE & SPEECH	ElevenLabs · Cartesia · Hume · Deepgram · AssemblyAI
----------------	--

VOICE AGENTS & TRANSPORT	Vapi · Retell · LiveKit · Pipecat · Twilio
--------------------------	--

GPU & CLOUD	AWS · GCP · Azure · CoreWeave · Modal · RunPod
-------------	--

VECTOR & DATA	Pinecone · Qdrant · Weaviate · Chroma
---------------	---------------------------------------

CODING TOOLS	Claude Code · Codex · Cursor · GitHub Copilot · Windsurf
--------------	--

How we connect: LiteLLM · Bifrost · Python SDK · JavaScript SDK · our open-source agent · provider admin APIs · CSV import · invoice ingestion

...and the rest of your stack. Where a vendor exposes usage, we pull it. Where they don't, we ingest the invoice.

A short call to your first spend map.

01

A short call

Your providers and your org chart. That's the whole agenda – no demo theater.

02

Your spend map

We come back with your spend picture built on your stack's shape and realistic numbers. Yours to keep either way.

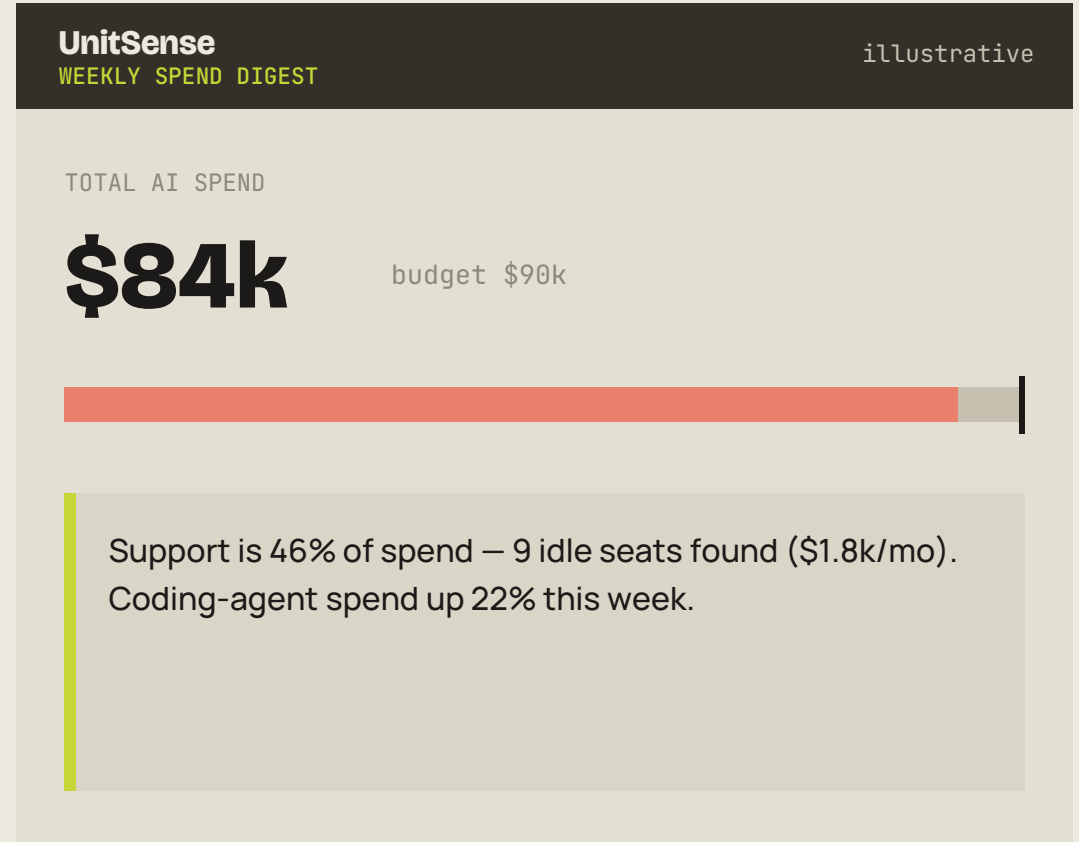
03

Built around you

If it's useful, we integrate against what you actually run. We work closely with a small number of early companies – the product gets shaped around your stack, not the other way around.

Your AI spend, in your inbox, every Monday.

No dashboard ritual. The digest lands with total spend vs. budget, what moved this week, and your single biggest saving — readable in 60 seconds, forwardable to your CFO.



We've run AI infrastructure at scale — now we're building its missing ledger.

Bilal Baqar

CO-FOUNDER & CEO

Director of AI Infrastructure at CTDS — built a petabyte-scale, zero-trust Kubernetes GPU platform. UChicago MS CS, Chicago Booth MBA. Previously PLUMgrid (enterprise SDN).

Salman Sikandar

CO-FOUNDER, PRESIDENT & CTO

Platform Engineering Lead at CTDS. Built and runs the UnitSense multi-tenant platform end to end. Previously PLUMgrid, Motorola.

Every AI dollar. Accounted for.

We'll show you what attribution and enforcement could look like — built on your providers and your org chart.

Setup a demo — unitsense.ai/book

more at unitsense.ai/governance