

AI ROI GUIDE

The AI Impact ROI Guide for VC & Investment Firms

.TrackImpact



CONTENTS

In This Guide

1 — Why This Matters	3
2 — The Maturity Ladder	5
3 — Self-Assessment	9
4 — ROI Framing	12
5 — Next Step	15

What ROI Means for Your Fund's Own Operations

Every fund now asks portfolio companies how they use AI — it's a standard diligence question, and increasingly a scoring input. Few funds turn that same lens on their own operations, or can put a number on what their own AI adoption is actually worth.



That's a real gap, not a rhetorical one. A fund's core product is judgment applied at speed: finding the right deals before competitors do, running diligence deep enough to catch what a founder's deck won't say, staying close enough to portfolio companies to catch trouble early, and turning all of it into LP communication that's accurate and on time. Every one of those functions is bottlenecked today by the same thing that bottlenecks a startup's own ops — partner and analyst hours spent on work that doesn't require partner-level judgment: reading data rooms for the first disqualifying signal, building comp tables, chasing portfolio-company updates that should already be flowing in, restating the same quarter's numbers three different ways for three different LP formats.

That work doesn't go away by wishing it away, and it doesn't get fixed by buying a tool. It gets fixed the same way it gets fixed inside a good portfolio company: by mapping the actual workflow, finding the parts that are boring and repeatable versus the parts that genuinely need a partner's judgment, building — or adopting — a system for the former so the humans spend their time on the latter, and then putting a number on what that shift is actually worth.

There's also a credibility angle a lot of funds haven't connected yet. A fund that can speak concretely about its own AI ROI — not "we use ChatGPT sometimes" but "here's where our diligence workflow sits, here's what we've changed, and here's what it's worth" — has a sharper, more specific answer than most peers when a founder, an LP, or a co-investor asks the question funds now ask everyone else. This guide, and the maturity ladder in the next section, is how a fund gets there.

Where AI Actually Delivers

✓ **List-Building & Screening**

✓ **Document Triage**

✓ **Time Back**

Where It Still Needs a Human

✗ **Partner-Level Judgment**

✗ **Founder & LP Relationships**

Where Does Your Fund Actually Sit Today?

A 6-rung scale, adapted for fund operations — Nobody, Person, Reactive-only, Cron, Looped, Agentic. The ladder isn't just a framework to read and nod at. It's a specific, repeatable process a fund can run this week to find out, function by function, where it genuinely sits.



The Process Behind Every Rung

Before a fund can honestly say where it sits on any rung, someone has to log the real, granular tasks that make up that function — not the department name, the actual task-level activity (not "diligence," but "read this data room and flag anything unusual," "build this comp table," "summarize this reference call"). Every logged task then gets scored for automation-readiness: does it genuinely need a partner's judgment, is it a good candidate for AI-assisted work with human review, or is it well-suited to running with minimal oversight?

That scoring is where most frameworks stop — and where the real question actually starts. A task scored as high-automation-potential isn't just "a task that could be automated once." The real question is whether it can become a standing loop — something that runs on a trigger, makes a bounded decision, takes an action, and gets better over time from its own results, without a human rewriting the process each cycle. That's the difference between doing a task with AI help once and building a system that keeps doing it and keeps improving. It's the actual mechanism behind every Agentic row below — not a vague aspiration, but the specific, repeatable outcome of tracking and scoring done consistently and acted on.

Deal Sourcing

Rung	What It Looks Like
Nobody	No systematic sourcing at all — deals arrive only through inbound referrals and whoever happens to be in the room.

Person	A partner or analyst manually scans Crunchbase/PitchBook, LinkedIn, and newsletters on an ad hoc basis, keeps a personal shortlist. Entirely dependent on that person's time and memory.
Reactive-only	A system exists but only responds to input — e.g. a shared CRM where deals get logged after someone finds them, or a saved search that has to be manually re-run. Nothing surfaces a deal on its own.
Cron	A scheduled pull — e.g. a weekly automated scrape/alert against a thesis (sector, stage, geography) lands in an inbox or Slack channel on a fixed cadence, same criteria every time, regardless of what did or didn't convert to a meeting last time.
Looped	The sourcing output gets reviewed and scored (a partner marks each surfaced deal pass/interesting/meeting), and that scoring measurably changes what gets surfaced next cycle — the source list, keyword set, or scoring weights get adjusted based on what actually turned into real meetings, not just re-run unchanged.
Agentic	Trigger → decision → bounded action → recorded outcome → next run uses that outcome, with no human writing the rules between cycles. A system watches funding announcements, hiring signals, and product-launch signal against the fund's stated thesis; independently decides which are worth a first-touch outreach; sends a bounded first message itself; logs whether it converted to a call; and adjusts what it prioritizes sourcing next based on that real conversion data — a human enters only when a founder actually replies.

Diligence

Rung	What It Looks Like
Nobody	No consistent diligence process — depth and rigor vary entirely by which partner is running the deal and how much time they have that week.
Person	A standard checklist exists, but every data room gets read manually front to back by an analyst, every reference call is scheduled and summarized by hand, every comp table is rebuilt from scratch each time.
Reactive-only	Tooling exists but only fires on manual trigger — e.g. someone pastes a data room link into an AI tool and asks it to summarize, on that one deal, that one time. No standing process, no memory across deals.
Cron	A standing pipeline runs the same automated checks on every deal that enters diligence — auto-extracting cap table structure, auto-flagging missing standard documents, auto-generating a first-pass comp set — on a fixed process every time, whether or not last quarter's version of that check actually caught anything useful.
Looped	Diligence-tool output gets checked against what the deal actually turned out to be, and that feedback changes the checklist or flagging thresholds — a category of red flag that turned out to be noise gets deprioritized; one that turned out to matter gets escalated in future runs.
Agentic	The system independently triages incoming data-room documents, flags the specific items that deviate from the fund's own historical pattern of "deals that later had problems" — not a generic checklist, a pattern learned from this fund's own outcome history — and routes only the genuine judgment calls to a partner, with the flagging logic itself updating from confirmed outcomes deal over deal.

Portfolio Monitoring

Rung	What It Looks Like
Nobody	No structured monitoring between board meetings — the fund finds out a portfolio company is in trouble when the company tells them, or when it's too late to help.

Person	A partner emails/calls each portfolio company periodically to ask "how's it going," takes notes manually, no shared system across the fund's whole portfolio.
Reactive-only	A system exists (e.g. a KPI-submission form or portal) but only produces signal when a company proactively fills it in — nothing chases a company that goes quiet, nothing compares this quarter's numbers to the pattern of companies that later struggled.
Cron	A scheduled check-in — a monthly auto-generated request for updated metrics goes out to every portfolio company on the same date regardless of that company's actual situation, and gets logged into a spreadsheet or dashboard without further action.
Looped	Submitted metrics get compared against expectation (this company's own trend, or the fund's benchmark for similar-stage companies), flagged when they deviate, and a partner's response to past flags changes what gets flagged next time.
Agentic	The system continuously ingests whatever signal is available (self-reported metrics, plus passively available signal like hiring-page changes or review-site sentiment where available), independently decides which portfolio companies show a genuine risk pattern versus normal noise, and escalates only the real cases to the responsible partner with a specific reason — with the risk model itself refined by which past escalations turned out to matter.

LP Reporting

Rung	What It Looks Like
Nobody	LPs get updates irregularly, format and content vary call to call, no standing cadence.
Person	Someone — often a junior partner or the fund's ops person — manually assembles the quarterly letter and financial summary from scratch every cycle, pulling numbers from several systems by hand, writing narrative fresh each time.
Reactive-only	A template exists, but every cycle still starts from a blank slate that has to be manually re-populated — no automated pull from the underlying data.
Cron	Numbers auto-populate into the report template on a fixed schedule (e.g. quarter-close triggers an automated pull of NAV, IRR, portfolio-company metrics into a draft), but the draft still needs full manual review and the process doesn't adapt to what LPs actually ask follow-up questions about.
Looped	LP questions and feedback from past reports — what they asked for clarification on, what they said was useful vs. noise — get tracked and change what the next report emphasizes or how it's formatted. The report gets measurably better at answering the questions LPs actually have, cycle over cycle.
Agentic	The system assembles a full first-draft LP report from underlying data sources on its own, flags anywhere a number looks anomalous or needs a partner's narrative judgment before it's said to LPs, and a partner's edits/approvals feed back into what the system drafts unprompted next cycle — while any specific claim or number is still gated through partner review before it reaches an LP.

Score yourself: for each function above, circle the rung that honestly matches your fund today — not where you'd like to be. Most funds land on different rungs for different functions, and that's the point: the self-assessment in the next section turns this into a specific, numbered starting point.

A Note on Where Funds Actually Cluster Today

Most funds are probably somewhere between Person and Cron across these four functions — sourcing and LP reporting are the two most likely to already have some Cron-grade tooling (a scheduled deal-flow alert, an auto-populated LP template), because those are the two functions with the most off-the-shelf software already built for them. Diligence and portfolio monitoring are more likely to still be Person-grade, because both require judgment that's harder to templatize and because the tooling market for them is younger.

This is a reasoned expectation based on how mature the tooling market is for each function — not a cited survey finding. If your fund has (or finds) a real fund-operations maturity study, use that instead of this paragraph.

12 Questions, 3 Per Function

The fastest way to answer these honestly isn't to guess — it's to actually log a week of real tasks across these four functions and look at what you find, using the same track-then-score process from Section 2. Answer what you can from memory today, then revisit after a week of real logging.



Deal Sourcing

1. If every partner at your fund were unavailable for a month, would deal flow visibly slow down — or does sourcing happen through a system, not just people?
2. Can you say, with a number, what fraction of last year's sourced deals came from a systematic process versus inbound/referral?
3. When a deal doesn't convert to a meeting, does that outcome change what your sourcing process looks for next time — or does the same process run unchanged regardless of results?

Diligence

4. Does every deal get the same baseline diligence checks, or does depth depend on which partner has time that week?
5. Is there a record of which diligence red flags, in your fund's own history, actually predicted a problem later — and does that record influence what gets flagged today?
6. How much analyst time does a typical diligence pass consume before it reaches a partner-level judgment call? (A rough hours estimate is fine.)

Portfolio Monitoring

7. If a portfolio company started struggling quietly between board meetings, how would your fund find out — and how fast?
8. Do you compare a portfolio company's current metrics against a pattern (its own trend, or similar-stage peers), or just look at the number in isolation each time?

9. Is there a standing process that flags a portfolio company for partner attention, or does that only happen when someone happens to notice?

LP Reporting

10. How many hours does producing one quarterly LP report take, start to finish, across everyone involved?

11. Does report content change based on what LPs actually ask about after past reports, or is the format fixed regardless of feedback?

12. If your ops/IR lead left tomorrow, could the next LP report still go out on time without them?

Rough Scoring Guidance

Mostly "no" / "nobody would know" answers across a function → that function is Nobody/Person-grade. Mostly "yes, there's a process, but it doesn't change based on outcomes" → Reactive-only/Cron. Any "yes, and we've changed our process based on what worked" → Looped or better. Very few funds will score Agentic on any function today — that's the point of this guide, not to imply failure, but to give a concrete target above wherever the fund currently sits.

TRACKIMPACT — SEARCH PROFILES

Not Sure Where to Start? Search Your Job Title.

TrackImpact's marketplace covers thousands of O*NET job titles, and every profile breaks the role down into its real, recurring tasks — not "use AI for marketing," but the actual activities someone in that seat does every week — each matched to a concrete AI approach and an automation-readiness read.

Instead of guessing which tasks are worth automating, search the job title and see the answers already mapped out.

SEARCH

"Senior Financial Analyst"

TASKS MAPPED

38 recurring tasks found

→ SAMPLE TASK MATCH

"Reconcile the variance report against last month, flag anything outside a 5% band."

HITL SCORE: MEDIUM — AI DRAFTS, YOU REVIEW

Search your job → www.trackimpact.io/search/

What Moving Up the Ladder Is Actually Worth

Reasoned conservatively, shown as arithmetic a reader can adjust for their own fund — not a single invented industry-wide number. Every figure below is either a clearly-labeled illustrative assumption or explicitly flagged as needing a real source before you'd want to state it externally.



Deal Sourcing

A small/mid-size fund's investment team plausibly spends 5-10 hours/week per person on manual sourcing research — scanning, list-building, first-pass screening. (Illustrative range, not a cited figure — swap in your own number if you track this.) Moving from Person/Reactive-only to Cron-grade automated surfacing plausibly reclaims the list-building portion of that, not the judgment portion — conservatively, 2-4 hours/week per person. At a fully-loaded analyst/associate hourly cost specific to your fund, that's real, recurring time. The honest claim is "hours back for higher-value work," not a dollar figure this guide can respectably state without a sourced cost basis.

Diligence

A full diligence pass on a single deal involves some number of hours of analyst time on first-pass, non-judgment work — document review, comp-table assembly, reference-call scheduling and note transcription. That figure varies enormously by stage and sector and shouldn't be guessed at with false precision. Even a conservative estimate that automation could reclaim a third of that first-pass work, per deal, compounds fast for a fund running multiple deals through diligence concurrently — the multiplier is deal volume, which you can plug in for your own fund without this guide needing to guess a global number.

Portfolio Monitoring

The clearest ROI case here is structural, not hours-based: time-to-detection on a portfolio company that's struggling. A fund that only finds out about trouble at the next scheduled board meeting has a detection lag of up to a full quarter; a Looped-or-better monitoring

process can shrink that to weeks. The economic value of catching a problem a quarter earlier is real but genuinely hard to conservatively estimate without a specific fund's own loss history or a published study correlating early detection with recovery outcomes — deliberately left unstated here rather than guessed at, since a reasoned-but-fake number would be worse than no number for something this high-stakes.

LP Reporting

This is the most measurable of the four, because report-production hours are usually already tracked informally even where nothing else is. If a quarterly LP report currently takes, conservatively, 15–25 hours of combined partner/ops/analyst time to assemble (illustrative range — varies by fund size and LP count), and moving from manual assembly to Cron-grade auto-population of the underlying numbers (leaving narrative and partner review as the human part) plausibly reclaims half of that, the arithmetic is straightforward for your own fund:

(current hours × 0.5) × 4 quarters × your fully-loaded hourly cost

The Real Roll-Up

The honest summary isn't "moving up the ladder saves \$X" — it's that every function above has a real, non-fabricated mechanism by which partner and analyst hours currently spent on non-judgment work get returned to judgment work. And that mechanism compounds specifically because of the loop structure described in Section 2: a task that's tracked, scored, and turned into a standing loop keeps paying off and keeps improving on its own — it's not a one-time hours-saved event, it's a system that gets better the longer it runs. A fund running this self-assessment gets a specific, function-by-function starting point for that conversation with its own numbers, which is worth more than a single invented industry-wide average would be.

TRACKIMPACT — TRACK & SCORE

Turn Everyday AI Use Into Evidence

Once a task is worth doing, TrackImpact is where it gets logged — the tool, the task, the time saved. Every entry is automatically scored for automation-readiness, and the highest-scoring tasks are flagged as candidates to become a standing loop: a process that runs on a trigger, takes a bounded action, and gets better from its own results.

Track, score, and step back — the whole loop, live today. That's how scattered AI usage becomes the evidence the next budget conversation runs on.

TOOL	TASK	TIME SAVED	SCORE
ChatGPT	Draft first-pass client email	15 min	82 — Loop candidate
Otter.ai	Transcribe client call	20 min	65
Excel Copilot	Reconcile month-end entries	35 min	40



↳ Recorded outcome feeds the next run

Get started → trackimpact.io

Where TrackImpact Fits

This is the same process TrackImpact runs on itself, right now — not a theoretical framework, but the actual loop the product is being built around.



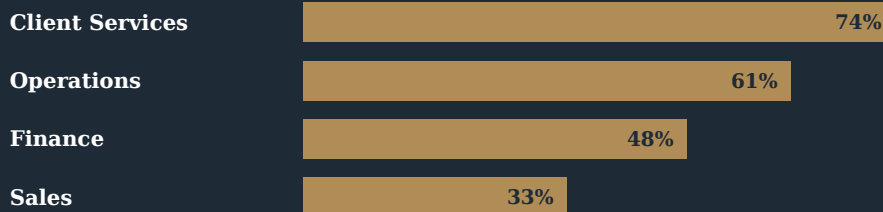
TrackImpact's own team runs this exact process on its own operations — logging real activity across its own functions, automatically scoring automation-readiness, and turning the highest-scoring tasks into standing loops. That's a stronger tie-in than a generic "try our product" close: this isn't a framework we're asking a fund to adopt on faith, it's the same process running inside the company that built the guide.

TRACKIMPACT — TEAMS

See Impact Across the Whole Organization

Teams brings every individual's tracked and scored activity together at the organizational level — so leadership can see adoption and value by department, site, or role, and compare who's driving results.

It's the difference between a handful of enthusiastic early adopters and a rollout that can actually be managed and reported on.



See how it works → trackimpact.io

Ready to Set Up TrackImpact for Your Fund?

Reach out and we'll get you started — mapping roles, setting up your team, sending onboarding invites to your staff, and walking through how the platform and reporting work.



Ade Molajo

Founder, TrackImpact

LinkedIn: [linkedin.com/in/ademolajo](https://www.linkedin.com/in/ademolajo)

Email: ade.molajo@trackimpact.io

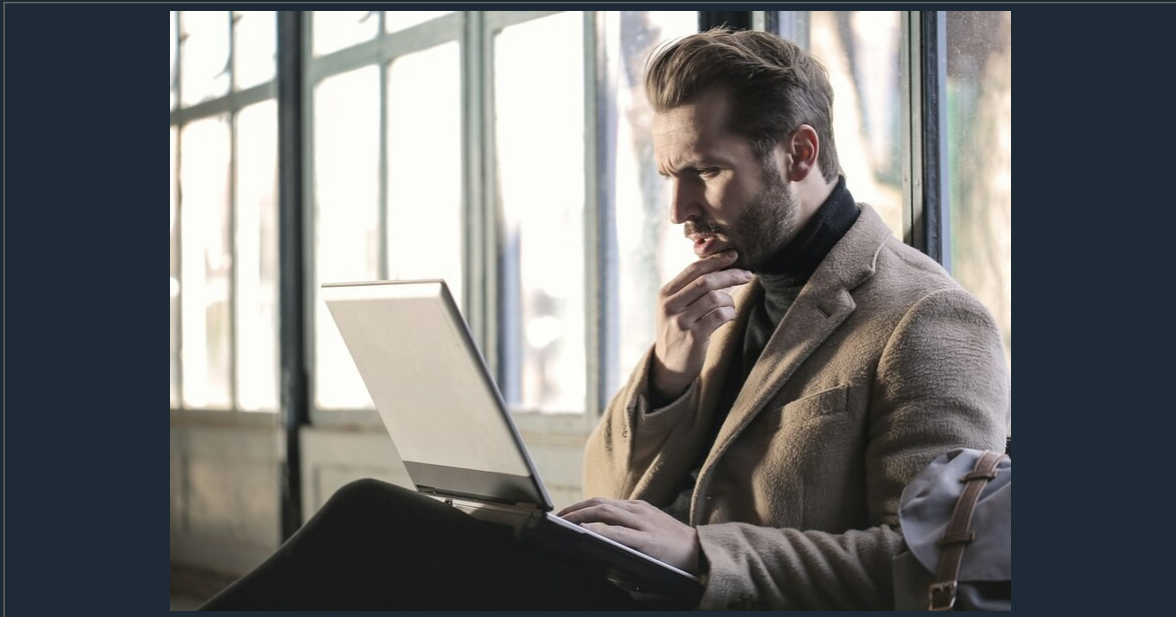
Web: trackimpact.io

READ NEXT

The AI Skills Guide

This guide made the case for AI at the organizational level. The AI Skills Guide takes it to the individual — the specific habits, prompts, and workflows that turn AI from a tool people dabble with into one they rely on daily.

A natural next step for any team lead or partner rolling this out person by person.



Read the AI Skills Guide → www.trackimpact.io/guides/ai-skills-guide/

Photo Credits

All photography sourced from Unsplash. The Unsplash License does not require attribution, but it's listed here as good practice.

- Team meeting in a modern office — Vitaly Gariev / Unsplash (img-vc-framing.jpg)
- Team collaborating around a whiteboard during a meeting — Vitaly Gariev / Unsplash (img-vc-ladder.jpg)
- Team collaborating around a whiteboard during a meeting — Vitaly Gariev / Unsplash (img-vc-selfassessment.jpg)
- Two colleagues reviewing documents together — Vitaly Gariev / Unsplash (img-vc-roi.jpg)
- Person at a desk with laptop and papers — Unsplash (img-vc-close.jpg)
- Statistics on a laptop — Carlos Muza / Unsplash (img-vc-cover.jpg)
- The AI Skills Guide cover — Unsplash (img-skills2-cover.jpg)