

AI ROI GUIDE

The AI Impact ROI Guide

A Playbook for Scaling AI Across Your Organization

.TrackImpact



CONTENTS

In This Guide

1 — Why This Matters	3
2 — The Maturity Ladder	5
3 — Self-Assessment	8
4 — ROI Framing	11
5 — Next Step	14

Adoption Isn't the Hard Part Anymore. Proof Is.

Almost every organization now has someone whose job is to make AI adoption real — not another pilot, not another tool announcement, but AI actually changing how work gets done at scale. Few of them can yet put a number on what that adoption is actually worth, function by function.



That gap is the job, not a side effect of it. A transformation lead or Chief AI Officer is measured on outcomes across functions they don't personally run day to day — sales and support, finance and ops, marketing and content, people and talent. Each of those teams may already have scattered AI use happening inside it. What's usually missing isn't more tools; it's a consistent way to see where AI is actually being used, whether it's working, and what it's worth — across the whole organization, not one team's dashboard at a time.

That visibility doesn't come from a survey or a slide asking teams to self-report their AI usage once a quarter. It comes the same way it does inside any function that's actually been systematized: by mapping the real, task-level activity, scoring it for automation-readiness, and tracking what happens next — consistently, across every function, not just the ones with an enthusiastic early adopter.

There's a credibility dimension too. A transformation lead who can say "here's exactly where AI is creating value across the org, here's what we've changed, and here's what it's worth" has a fundamentally stronger position — with the board, with budget owners, with skeptical function heads — than one relying on anecdotes and adoption-rate vanity metrics. This guide, and the maturity ladder in the next section, is how that position gets built.

Where AI Actually Delivers

✓ Cross-Functional Visibility

✓ Task-Level Tracking

✓ Time Back

Where It Still Needs a Human

✕ **Function-Head Judgment**

✕ **Org-Specific Risk Calls**

Where Does Each Function Actually Sit Today?

A 6-rung scale — Nobody, Person, Reactive-only, Cron, Looped, Agentic — applied across four functions every organization has. It isn't a framework to read and nod at. It's a specific, repeatable process a transformation lead can run this week to find out, function by function, where the organization genuinely sits.



The Process Behind Every Rung

Before anyone can honestly say where a function sits on any rung, someone has to log the real, granular tasks that make it up — not the department name, the actual task-level activity (not "customer support," but "draft a first reply to this ticket type," "summarize this call," "flag this account as at-risk"). Every logged task then gets scored for automation-readiness: does it genuinely need a human'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 a team trying AI once and an organization running a system that keeps improving. It's the actual mechanism behind every Agentic row below.

Sales & Customer Success

Rung	What It Looks Like
Nobody	No AI use in the sales or support workflow at all — every email, every reply, every follow-up written from scratch by a person.

Person	A rep or support agent personally pastes things into ChatGPT or Claude ad hoc — drafting an email, summarizing a call — entirely dependent on that individual remembering to use it.
Reactive-only	A tool exists (e.g. an AI drafting assistant in the CRM or helpdesk) but only fires when someone manually triggers it on a single ticket or email. Nothing runs on its own.
Cron	A standing process — every incoming ticket gets an AI-drafted first-pass reply queued automatically, or every call gets an AI summary attached to the CRM record — runs the same way every time, regardless of outcome.
Looped	Which AI-drafted replies went out as-is versus got heavily rewritten, and which AI call summaries actually got used in the next call, gets tracked — and that feedback changes what the system drafts next time.
Agentic	Trigger → decision → bounded action → recorded outcome, with no human writing the rules between cycles. The system watches inbound tickets, decides which match a known resolved pattern, sends a bounded first response itself, logs whether it resolved without escalation, and adjusts what it auto-handles next based on that real data — a person enters only when something needs judgment a policy doesn't cover.

Finance & Operations

Rung	What It Looks Like
Nobody	No AI use in reporting or operations — every reconciliation, every variance report, every ops metric compiled entirely by hand.
Person	Someone personally pastes numbers into an AI tool occasionally to help write a summary or sanity-check a formula, informally, with no consistent workflow.
Reactive-only	A tool exists (e.g. Excel Copilot) but only gets used when someone remembers to open it on a specific spreadsheet that week.
Cron	A scheduled process — month-end variance flags auto-generated on a fixed date, or an ops summary auto-compiled every Monday — runs the same way every cycle whether or not last cycle's version caught anything useful.
Looped	Which flagged variances turned out to be real issues versus noise gets tracked, and that feedback changes what gets flagged, or how sensitive the thresholds are, next cycle.
Agentic	The system continuously reconciles incoming financial data against expected patterns, independently decides which variances are genuine issues versus rounding noise, drafts the exception note itself, and logs which of its past flags a controller actually confirmed — refining its own thresholds from that outcome data, with any number that reaches a stakeholder still gated through human review.

Marketing & Content

Rung	What It Looks Like
Nobody	No AI in the content or campaign workflow — every post, every ad variant, every brief written entirely from scratch.
Person	A marketer personally uses an AI tool ad hoc for a first draft of a post or email — informally, tool by tool, with no shared process across the team.
Reactive-only	A content-AI tool exists in the stack but only produces output when someone manually opens it and asks for a specific piece. Nothing runs unprompted.
Cron	A scheduled process — a first-draft content calendar auto-generated weekly from a content plan, or ad-copy variants auto-generated for every new campaign — runs the same way every time.

Looped	Which AI-drafted variants actually got published versus rewritten, and which performed well versus poorly once live, gets tracked — and that feedback changes what the system proposes next cycle.
Agentic	The system monitors campaign performance, independently decides which underperforming variants to retire and which angles to generate more of, drafts and queues the new variants itself, and logs which of its own past calls actually improved performance — adjusting its own generation strategy from that data, with brand-sensitive copy still routed through a marketer's review before it goes live.

People & Talent

Rung	What It Looks Like
Nobody	No AI use in hiring or people ops — every resume screened, every job description written, every onboarding step handled entirely manually.
Person	A recruiter or HR generalist personally uses an AI tool ad hoc to draft a job posting or summarize a resume — individually, with no shared team process.
Reactive-only	An AI screening or drafting tool exists in the ATS but only fires when someone manually runs it on a specific requisition.
Cron	A scheduled process — every new applicant gets an AI-generated first-pass summary automatically, or a standard onboarding checklist auto-populates for every new hire — runs the same way every time.
Looped	Which AI-flagged candidates actually made it to interview versus got correctly screened out gets tracked, and that feedback changes screening criteria or emphasis next cycle.
Agentic	The system independently triages incoming applications against the role's actual hiring pattern — not just keyword match, a pattern learned from which past hires actually succeeded — decides which candidates to fast-track to a recruiter, and logs which of its recommendations led to an actual hire, refining its own screening model from that outcome data. Every real hiring decision is still made by a person.

Score it: for each function above, circle the rung that honestly matches your organization today — not where you'd like to be. Most organizations 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 for each one.

A Note on Where Organizations Actually Cluster Today

Most organizations are probably somewhere between Person and Cron across these four functions — marketing and sales are the two most likely to already have some Cron-grade tooling, because those are the two functions with the most off-the-shelf AI software already built for them. Finance and people functions are more likely to still be Person-grade, because both involve more judgment that's harder to templatize and 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 organization has (or finds) a real cross-functional AI 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.



Sales & Customer Success

1. If every rep or agent were out for a week, would response times to inbound leads or tickets visibly slip — or does a system carry the load?
2. Can you say, with a number, what fraction of first-response drafts across sales and support are AI-assisted today?
3. When an AI-drafted reply gets heavily rewritten by a human, does that change what the system drafts next time — or does it keep making the same kind of miss?

Finance & Operations

4. Does every recurring report get the same baseline AI-assisted first pass, or does it depend on who has time that week?
5. Is there a record of which AI-flagged variances turned out to be real issues versus noise — and does that record change what gets flagged today?
6. How many hours does a typical month-end or recurring reporting cycle consume before it reaches a human judgment call? (A rough estimate is fine.)

Marketing & Content

7. Could your content calendar keep moving for a week without a person manually generating every first draft?
8. Do you compare AI-drafted content's actual performance against human-drafted content, or just assume one is fine?

9. Does a piece of content's real-world performance change what your AI tools generate next, or does the prompt stay the same regardless of results?

People & Talent

10. If your recruiter or HR generalist left tomorrow, would new-hire onboarding still run on schedule without them personally driving each step?

11. Is candidate screening based on a pattern of who actually succeeded after being hired, or just keyword and resume matching in isolation?

12. How much of a typical hiring cycle is spent on work that doesn't require a human judgment call yet — and is any of it currently automated?

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 organizations 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 each function 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 organization — 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.



Sales & Customer Success

A sales or support team plausibly spends a meaningful share of rep/agent time on first-draft replies and call summaries — work that doesn't require judgment, just needs to get done consistently. Moving from Person/Reactive-only to Cron-grade automated drafting plausibly reclaims the drafting portion of that, not the judgment portion — conservatively, a couple of hours per person per week. At a fully-loaded rep or agent hourly cost specific to your organization, that's real, recurring time redirected toward the parts of the job that actually need a person: complex accounts, escalations, relationship-building.

Finance & Operations

A recurring reporting cycle involves some number of hours of first-pass, non-judgment work — data pulls, variance flagging, formatting the same numbers into different templates. That figure varies enormously by org size 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 cycle, compounds fast for an organization running the same cycle monthly or quarterly across multiple teams — the multiplier is cycle frequency and team count, which you can plug in for your own organization without this guide needing to guess a global number.

Marketing & Content

Content production is one of the more measurable cases, because output volume is usually already tracked. If a marketing team currently spends, conservatively, several hours per piece on first-draft generation across a given volume of content per month (illustrative range

— varies hugely by team size and content mix), and moving from manual drafting to Cron-grade AI-assisted first drafts plausibly reclaims a meaningful share of that — leaving strategy, brand judgment, and final review as the human part — the arithmetic is straightforward for your own team:

(pieces per month × current hours per piece × reclaimed %) × fully-loaded hourly cost

People & Talent

The clearest structural ROI here isn't hours-based, it's time-to-fill and screening quality: a team relying entirely on manual resume review has a hard ceiling on how many requisitions it can meaningfully triage at once. Looped-or-better screening — where the system learns from which candidates actually succeeded, not just keyword match — plausibly shrinks both the time spent per requisition and the rate of good candidates missed. The economic value of a faster, better hire is real but genuinely hard to conservatively estimate without your organization's own cost-per-hire and time-to-fill data — deliberately left unstated here rather than guessed at, since a reasoned-but-fake number would be worse than no number for something this consequential.

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 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 transformation lead running this self-assessment across every function gets a specific, function-by-function starting point for that conversation with their own numbers, which is worth more than a single invented org-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.



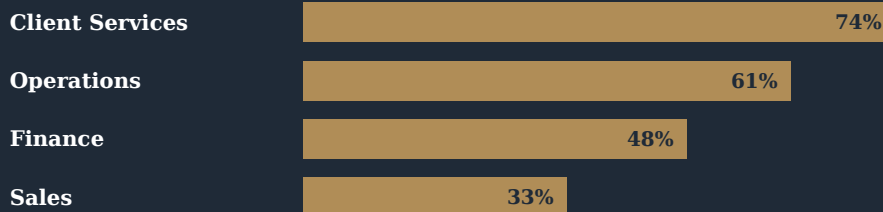
TrackImpact's own team runs this exact process on its own operations — logging real activity across its own functions, automatically scoring automation-readiness, and turning the highest-scoring tasks into standing loops. That's a stronger tie-in than a generic "try our product" close: this isn't a framework we're asking an organization 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

This guide is the organizational half of the picture. The AI Skills Guide is the individual half — the six habits that make someone genuinely effective with AI day to day, which is what actually moves the needle once the tracking and reporting from this guide are in place.

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

Ready to Set Up TrackImpact for Your Organization?

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



Ade Molajo

Founder, TrackImpact

LinkedIn: 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 transformation lead rolling this out team by team, 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-generic-framing.jpg)
- Team collaborating around a whiteboard during a meeting — Vitaly Gariev / Unsplash (img-generic-ladder.jpg)
- Team collaborating around a whiteboard during a meeting — Vitaly Gariev / Unsplash (img-generic-selfassessment.jpg)
- Two colleagues reviewing documents together — Vitaly Gariev / Unsplash (img-generic-roi.jpg)
- Person at a desk with laptop and papers — Unsplash (img-generic-close.jpg)
- Team gathered around a laptop, reviewing something together — Vitaly Gariev / Unsplash (img-generic-cover.jpg)
- The AI Skills Guide cover — Unsplash (img-skills2-cover.jpg)