



AcceleratorApp
Software for Accelerator & Incubators

2026 EDITION

The AI Application Scoring Playbook

Review every applicant well. Decide every applicant yourself.

A practical guide to setting up, running, and getting the most from AcceleratorApp's AI evaluator across application rounds of any size.

Every program that grows eventually hits the same ceiling.

Applications increase. Evaluator hours do not. The response is always one of three things: review everyone shallowly, miss the decision deadline, or recruit ad-hoc reviewers whose scoring is all over the place. None of these options improves selection quality. They manage a constraint.

The constraint is human hours applied to the wrong work. Sorting the clearly-strong from the clearly-weak is mechanical: it requires reading against a set of criteria and producing a score. It does not require the judgment your best evaluators bring to the genuinely hard calls in the middle of the pool.

"The AI evaluator handles the mechanical part. Your team handles the judgment."

WHAT THIS GUIDE COVERS

Writing a thesis the AI can use well. Setting up evaluation form instructions. Running the evaluator and reading its output. Using the rejection feedback workflow. Measuring quality over time. What to do when something looks wrong.

The AI evaluator is an evaluator type you assign to a round, exactly like a human reviewer.

Once assigned, it reads each application against your funnel's thesis and the instructions you attach to the evaluation form. It scores every question, produces per-question reasoning, and ranks the pool. Your human reviewers open a shortlist with rationale attached.

What it reads

- Your funnel thesis
- Your evaluation form questions and criteria
- Your form instructions
- The applicant's answers and attachments

What it produces

- A score per question with reasoning
- A ranked application pool
- A tailored rejection feedback draft

What it does not do

- Make a final decision
- Contact applicants
- Score beyond your criteria

INPUTS DECIDE EVERYTHING

The AI is as good as the inputs you give it. A vague thesis produces vague scores. A well-written thesis with specific signals, anti-signals, and calibration anchors produces evaluations your team can trust and act on.

The thesis is the single most important input to every AI evaluation.

It tells the AI what a strong applicant looks like for your specific program. Not a generic notion of a promising startup. Your definition, for your funnel.

- F** **Focus.** Who this funnel is for: sector, stage, geography, founder profile. Be narrow. A thesis that excludes nothing cannot discriminate.
- I** **Indicators.** What to look for and what to weigh against. Both positive signals and anti-signals are required. Positive-only theses produce overly generous scores.
- T** **Thresholds.** What excellent, acceptable, and weak look like for this funnel. Calibration anchors prevent scoring from drifting away from your team's.

THESIS WRITING CHECKLIST

- ☐ Names the sector and stage specifically, not generically
- ☐ Lists at least two concrete positive signals and two anti-signals
- ☐ Includes calibration anchors: what excellent, acceptable, and weak look like
- ☐ Avoids marketing language: innovative, transformative, world-class, scalable
- ☐ Is 200 to 2,000 characters, and scored 7 or above by the Prompt Coach

REQUIRED BEFORE ASSIGNMENT

Run the Prompt Coach before assigning the evaluator. The Coach scores your thesis against the FIT criteria and surfaces specific improvements. It exists because the most common failure mode is a vague thesis producing vague evaluations.

Form instructions tell the AI how to evaluate answers to each question on your evaluation form.

They complement the thesis: the thesis defines good fit, the form instructions define how to read the specific answers on this form.

W

Weight. Which questions matter most. Without weighting, every question gets equal pull on the overall score, which rarely matches reality.

A

Attention. What to look for in each answer. The instruction tells the AI what a good answer to an open-ended question looks like.

G

Gradient. What the numeric scale means on this form, so a 7/10 from the AI means the same as a 7/10 from a human.

E

Edge cases. How to handle missing answers, partial answers, or broken links.

FORM INSTRUCTIONS CHECKLIST

- ☐ Identifies the one to two questions that matter most
- ☐ Notes what a strong answer looks like for each key question
- ☐ Includes scale calibration: what 10, 7, 5, 3, and 1 mean on this form
- ☐ Handles missing answers and broken links explicitly
- ☐ Does not repeat the thesis (the AI already has it as a separate input)

THE MOST SKIPPED STEP

Scale calibration is the most skipped step. Without it, AI scores and human scores mean different things, and your comparison view becomes unreliable. Two sentences on what each score band means is all it takes.

The AI evaluator slots into the same round assignment you already use for human reviewers.

ASSIGNMENT CHECKLIST

- ☐ Thesis is written and scored at 7 or above by the Prompt Coach
- ☐ Form instructions are written for each evaluation form in this round
- ☐ Form instructions are scored at 7 or above by the Prompt Coach
- ☐ Feature flag is enabled for this incubator
- ☐ Open round settings and assign the AI evaluator alongside (not instead of) human reviewers
- ☐ Confirm the "AI evaluator assigned" confirmation appears, showing how many forms have instructions

ALONGSIDE, NOT INSTEAD

Assign alongside humans, not instead of them. The AI evaluator works alongside your human panel. A round still requires a human decision to close. The AI surfaces the shortlist; your team makes the calls.

Your reviewers open the application overview and find a ranked pool with AI scores and reasoning attached.

#	APPLICANT	SCORE	LEAD REASONING
01	Helio Materials	9.2	Strong technical moat, named pilot customers, clear stage fit.
02	Tidal Health	8.4	Strong team and traction; regulatory path under-explained.
03	Loom Logistics	7.1	Good market, thin evidence of differentiation vs incumbents.
04	Ferro Robotics	5.8	Out of stage focus; early prototype, no validation.
05	Verde Pay	3.9	Anti-signal: crowded space, no clear wedge. Broken demo link.

READING AI SCORES EFFECTIVELY

- ☐ Start with the per-question reasoning, not just the overall score
- ☐ Treat a high score with weak reasoning as a flag; a low score with compelling reasoning as worth a manual look
- ☐ Watch the distribution: if everything scores 7 to 8, the thesis may lack discriminating power
- ☐ Use Re-run Coach if the distribution looks wrong before the review window closes

WHERE THE VALUE SHOWS

The comparison view is where the value shows. Reviewers seeing ten applications side by side with AI reasoning attached make faster, better-calibrated decisions than reviewers seeing ten applications cold.

For every applicant not advancing, the AI drafts tailored feedback grounded in their specific application.

Your team reviews the draft and sends it. Every founder learns why, instead of receiving the same generic form email.

REJECTION FEEDBACK CHECKLIST

- ☐ Review the AI draft before sending, not after
- ☐ Edit for tone: the draft is accurate but may need warming
- ☐ Confirm it references the specific criteria the application did not meet, not generic reasons
- ☐ Send from a named person, not a no-reply address
- ☐ Do not send AI-drafted feedback verbatim until you have calibrated the tone over the first few rounds

REJECTION QUALITY COMPOUNDS

In a small ecosystem, rejection quality compounds. Founders who receive specific, respectful feedback are more likely to reapply stronger, more likely to refer strong applicants, and less likely to tell peers your program is not worth applying to.

After your first two to three cohorts, you will have enough data to see whether the AI evaluator is calibrated well.

QUALITY INDICATORS TO WATCH

- ☐ **Staff acceptance rate.** What percentage of AI-scored applications are kept by your team? Target 70% or above in steady state.
- ☐ **Score distribution.** Is the pool ranking in a reasonable curve, or clustered at one end? Clustering suggests weak thesis calibration.
- ☐ **Thesis score correlation.** Are the high-scoring applications the ones your team advanced? Growing alignment signals a well-calibrated thesis.
- ☐ **Rejection feedback acceptance.** Sending drafts as-is or rewriting heavily? Heavy rewrites mean the tone needs adjusting at the form-instruction level.

FLAT DISTRIBUTION, FIX THE THESIS

If score distribution looks flat, re-run the Prompt Coach on your thesis. Flat distributions (every application scoring 6 to 8) almost always indicate a thesis with insufficient anti-signals or missing calibration anchors.

The AI evaluator is reliable but not infallible. Here are the common issues and what to do.

Scores look too generous or too harsh

Re-run the Prompt Coach on the thesis. Distribution problems almost always trace back to the thesis, not the model.

Per-question reasoning does not match your team's read

Check the form instructions. The attention guidance in the WAGE framework is probably missing or unclear.

A run failed or is showing an error

A retry button appears on the failed application. Retry consumes a new quota unit. Failed runs do not count against your monthly quota.

The evaluator is paused on a funnel

The thesis or form instructions were edited after scoring and the score is now stale. Re-run the Prompt Coach to clear the pause.

Something is systematically wrong with the reasoning

The evaluator snapshots the thesis, form instructions, and model version per run. Your CS contact can pull the snapshot to audit what context it used.

AcceleratorApp is a program management platform trusted by 200+ organizations across 6 continents, including German Accelerator, SparkLabs, CGIAR, MIT Design X, and Yale.

It runs the full program lifecycle in one system: application intake, AI application scoring, AI applicant coaching, multi-stage review, mentor matching, cohort management, milestone tracking, events, community, AI mentor intelligence, AI reporting, and AI translation.

HIGHLIGHTS

2,000,000+ applications processed

6,000+ program managers

ISO 27001 certified

GDPR compliant

Advisory-only AI, human decision required

Live in 1 to 2 weeks

Plans from \$499 per month

GET STARTED

Book a demo: acceleratorapp.co/en/demo

See pricing: acceleratorapp.co/en/pricing

Contact: support@acceleratorapp.co

ONE CONNECTED SYSTEM

- Application intake with AI application scoring
- AI applicant coaching before submission
- Multi-stage review and mentor matching
- Cohort management and milestone tracking
- Events and community
- AI mentor intelligence and AI reporting
- AI translation across the whole platform

The AI evaluator is advisory only. It scores, ranks, and drafts; a human decision is always required before a round closes.



AcceleratorApp
Software for Accelerator & Incubators

Ready to score every application against your thesis?

A 30-minute demo tailored to your application rounds, review panel, and selection criteria.

Book a demo · acceleratorapp.co/demo

See pricing · acceleratorapp.co/pricing

acceleratorapp.co · ISO 27001 certified · Used by 200+ programs globally