

Intelligence for the physical world.

AR without glasses. Intelligent light that guides the next action.

The Light Company is building augmented reality that appears directly on the world around us. Projected light shows people what to do, where to do it, and what comes next. Cameras connect that guidance to the real task. The vision is a system that understands physical work and guides it step by step, from the first action to completion.

Make the world the interface

A digital instruction still leaves a person to translate it into a physical action. Which part? Which position? Which direction? Prism brings the instruction to the point of action. An outline can show where an object belongs. A highlight can identify the next component. Verification can tell the system when to advance.

This is the company's original vision for AR: shared, useful guidance on real surfaces, without a headset or glasses. Robotics evaluation is the first commercial wedge into that larger opportunity.

10x

EVALUATION EFFICIENCY

By run count: 300 to 30

92.5%

LOWER REPORTED VARIABILITY

40% to 3%: relative reduction

10 units sold and deployed. 3 customers. \$15,000 in hardware sales value.
Enact 5 | New Theory 4 | Trossen Robotics 1. Founder-reported results and sales. [1, 2]

Replace guesswork with measured engineering

A robotics team needs to know whether its new policy improved. If every physical reset changes the test, that decision becomes guesswork. Prism makes the starting condition repeatable and verifiable. Teams can compare like with like, investigate failures, and build the next generation of physical AI through measurement.

The first product makes physical tests repeatable. The larger vision makes physical work intelligently guided.

The test changes before the robot moves.

Software can copy a test input. A robot lab has to recreate part of that input in the physical world. That is the initial-condition problem.

A fixed station can still produce a different test

The robot, table, cameras, and lighting can remain fixed. After every attempt, someone still has to put the objects back. Position, orientation, spacing, and arrangement become part of the next input. If those conditions change unintentionally, the team is no longer making the same comparison.

Imagine comparing two robot policies on a picking task. The first sees an object near the center of the table. The second sees it slightly rotated or closer to the edge. A different result can reflect the policy, the placement, or both. Without a reliable starting state, the team has less confidence in what the result means.

The cost spreads through the learning process

Evaluation informs which model to keep, which failure to investigate, and which data to collect next. Uncontrolled setup variation can send each of those decisions in the wrong direction. A false regression can trigger unnecessary debugging; a false improvement can hide a weakness. Initial conditions are therefore part of the measurement system, not a minor setup detail. Research on robot evaluation identifies this need for explicit conditions and richer evidence. [4]

What Prism solves

Prism turns a desired placement into a physical instruction. The operator resets the scene using projected guidance. The calibrated camera system checks the placement against the target and its tolerance. Guidance is removed before the robot acts, and the test result can be linked to its starting state. [2]

DEFINE	PROJECT	VERIFY	BLANK	RUN	RECORD
Target state	Show placement	Check the reset	Remove the cue	Test the policy	State + outcome

Define the physical input. Recreate it. Verify it. Then compare the policy.

This creates two useful modes: repeat the same states to compare policy versions, or vary states deliberately to find where a policy fails. Prism addresses reset consistency; robot behavior and sensor uncertainty still require repeated trials and sound measurement.

A concrete need. Three customers.

Customer	Deployed	Commercial signal
Enact (YC S26)	5	Needed one unit; bought five after one demo to improve initial-condition setup.
New Theory	4	Four units in the field for robot evaluation.
Trossen Robotics	1	One unit in the field; pilot underway.
Total	10	\$15,000 hardware sales value at \$1,500 per unit.

All ten units are deployed in the field, according to the founder. Hardware sales value reflects ten units at \$1,500 each; it is not a cash-collection figure. [1]

Why Enact bought five after one demo

The founder reports that Enact needed consistent initial-condition setup. Operators received placement guidance directly in the scene, helping them repeat the intended setup rather than reconstruct it from memory. Enact initially needed one unit and purchased five after the demonstration. [1]

The reported evaluation workload fell from 300-400 runs to about 30. Using the conservative 300-run baseline, that is a **10x improvement in evaluation efficiency by run count: 90% fewer runs**. The customer rationale is direct: more consistent setup reduces the need to repeat noisy evaluations.

10x

EVALUATION EFFICIENCY

By run count: 300 to 30

92.5%

LOWER REPORTED VARIABILITY

40% to 3%: relative reduction

The variability figure is the reported 40% to 3% result associated with New Theory's earlier evaluation workflows. The 300-to-30 run-count result is the founder's latest Enact account. These are separate reported outcomes, not a combined controlled study. Equal statistical confidence and end-to-end speedup have not been independently established. [1, 2]

Enact's published workflow also includes recreating failure states and collecting recovery data. This explains why repeatable physical inputs fit its daily work. YC lists Enact in its Summer 2026 batch. [3]

Find the missing data, not just more data

A large dataset can still leave some positions, orientations, or situations undersampled. The planned capability maps would show where tests succeed, where they fail, and where coverage is thin. Teams could target those gaps, collect relevant demonstrations, retrain, and repeat the same test. A failure map identifies where to investigate; comparison with training coverage helps determine whether missing data is the cause.

One hardware sale. Repeated software use.

The current hardware price is \$1,500 per unit. Software pricing is being explored as a subscription, potentially around \$100 per month, or usage-based billing. As customers train, test, and collect more data through Prism, paid activity can grow. Pricing and packaging remain open. [1]

Price around the work being completed

In robot evaluation, one candidate billing unit is a completed, verified reset recorded for a trial. In other guided workflows, the unit could be a completed task. The final unit and rate should match customer value and be easy to audit. Failed checks and repeated camera observations would not automatically become separate billable events.

Illustrative annual usage revenue

The following model assumes **\$1 per verified reset**, **100 billable resets per station per day**, and **250 operating days per year**. These are scenario assumptions, not current pricing, measured utilization, or a forecast. Hardware revenue is excluded.

Active paying stations	Annual verified resets	Annual usage revenue
10	250,000	\$250,000
100	2,500,000	\$2,500,000
1,000	25,000,000	\$25,000,000

Stations x useful activity x price per use

Growth can come from more customers, more stations per customer, and more work per station.

At 100 stations, halving either usage or price reduces this scenario to \$1.25 million annually; halving both reduces it to \$625,000. The economics depend on paid adoption and sustained use, not simply on units shipped.

Why the architecture offers a route to scale

Prism adds guidance and sensing around an existing workcell. An operator performs the physical reset, avoiding a dedicated reset robot for every task. A repeatable station design and reusable software can spread across more workcells. Calibration, installation, support, and hardware replacement remain the costs that must be standardized.

The \$1,500 selling price is established; hardware margin is not. The business case depends on recurring software revenue exceeding station delivery and support costs.

Build the wedge. Expand the vision.

A focused route to a much larger product

First: make robot tests repeatable. Deliver dependable guidance and verification, document reset accuracy, and establish the customer's baseline. Convert early use into regular paid activity.

Next: expose failures and coverage gaps. Map performance by physical condition and track how often each condition was tested. Compare these maps with available training coverage, collect data for undersampled states, and retest after training. Measure whether the failure region shrinks.

Then: extend intelligent guidance into physical work. Assembly, repair, training, and other tasks use the same interaction: understand the current state, show the next action, check the result, and continue. Each new workflow requires its own task knowledge and verification.

What can compound

Reusable workflows, calibration tools, integrations, and customer task history can make the system more valuable with repeated use. The goal is software that knows how to guide a task and recognize its completion. Customer data remains subject to agreed rights; cross-customer access is not assumed.

Alternative approaches include fixtures, internal procedures, and autonomous reset systems such as AutoEval. [5] Prism's position is flexible projected guidance with state verification in the customer's workspace. Its advantage must come from dependable execution and useful software, not projection alone.

Estimated raise: \$1.5 million

To standardize stations, improve reliability, build usage-based software, and expand paid customer workflows.

Notes and sources

[1] Founder update: all 10 sold units are deployed across Enact (5), New Theory (4), and Trossen Robotics (1), at \$1,500 each; Trossen pilot underway. Enact bought five after one demo for initial-condition setup; reported runs fell from 300-400 to 30. A subscription around \$100/month or usage-based pricing is under consideration.

[2] Supplied August final deck, pp. 4-5 and 14, and full explainer, pp. 4-7: product sequence and New Theory field report. $300 / 30 = 10$; $(40 - 3) / 40 = 92.5\%$. Metric definitions and trial logs remain outstanding. The raise estimate carries forward the final deck.

[3] Y Combinator: [Enact](#) and [Enact company website](#): company activity and workflow, checked for this revision. Public descriptions establish workflow fit; purchase rationale comes from the founder.

[4] Kress-Gazit et al. (2024), [Robot Learning as an Empirical Science](#).

[5] Zhou et al. (2025), [AutoEval](#), CoRL / PMLR 305. Research provides context, not validation of Prism's reported results.

Jonathan Solomon, Founder / CEO | hi@lightcompany.ai | lightcompany.ai