

Prism

Reliable robot evaluations. Guided by light.

We help physical AI labs find out
whether their models actually improved.

Our first step toward AR without glasses.



Prism in a real robot workcell

The team behind The Light Company

We built Prism beside lab operators, bringing repeatable physical tests into the customer’s workspace.



Jon Solomon
Founder

 Apple  HeyClicky

Helped ship Apple Vision Pro.
Led two stability labs with 150 robots and fixtures.



Pranav Takrani
Engineering team

  Northstar Robotics



Steph Ng
Engineering team

 Facebook AR/VR
SONY **USC**

Computer vision, robotics and AR/VR.



Jon Cosgrove
Advisor

 Apple

Lead Robotics @ Apple.

Did the model improve, or did the setup change?

A lab changes its robot's model.
An operator resets the cup.
The cup lands in a different place.

The next result now reflects both changes.
Teams repeat trials to work out whether the
model got better.

**Identical test code does not
create identical physical conditions.**



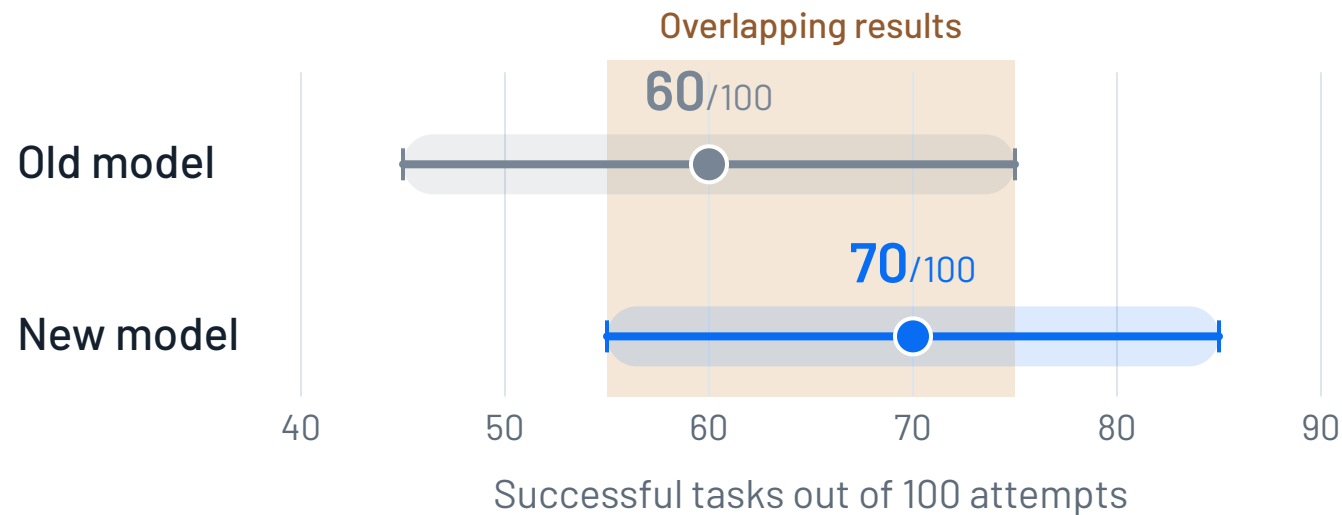
Position and orientation change what the robot must do.

Make the model improvement visible.

Object placement changes the task. Prism makes the starting conditions repeatable.

Before Prism

Object positions change between resets.

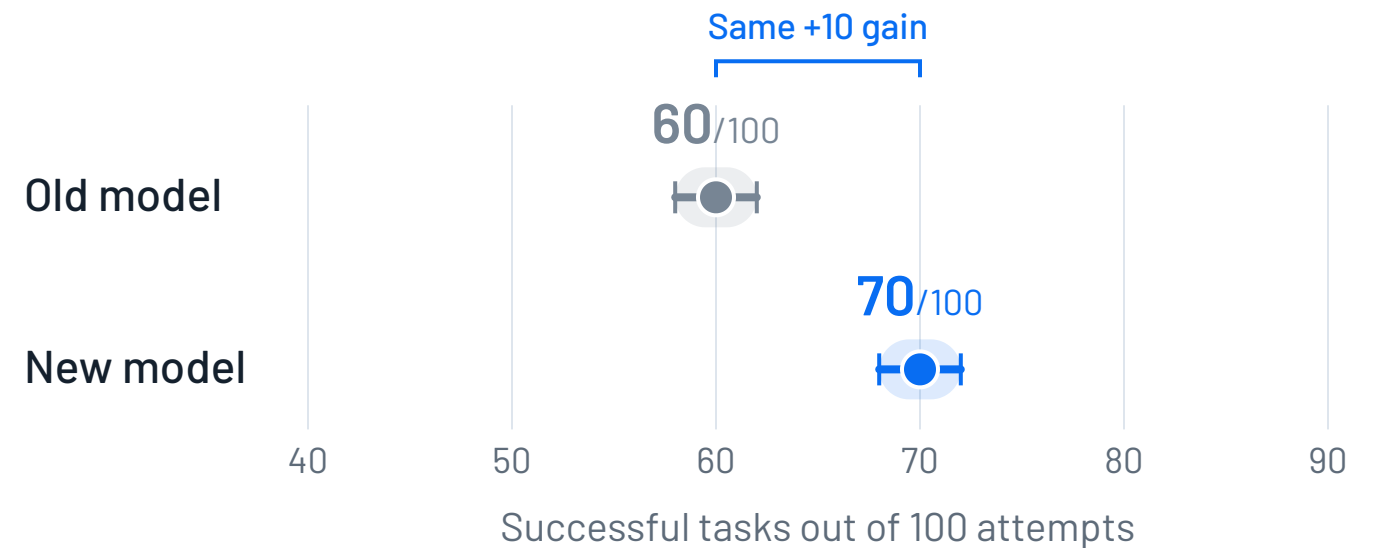


Variation hides which model is better.

Illustration · Same two models. Same averages. Same scale.

With Prism

Light guides the placement. Cameras verify it.



The same model improvement stands out.

● Average — Spread across repeated tests

New Theory

Measured over 30 days

Before
40% → 3%
With Prism

**92.5% less
evaluation variability.**

Engineers can see whether their model improved.

Charts illustrate the mechanism; they are not customer measurements. New Theory's relative reduction = $(40 - 3) \div 40$.

Prism connects the test to the real scene



01 / Guide the placement

Upload or generate a case. Prism projects where the object belongs.



02 / Verify the conditions

The operator resets the object.
Cameras check the physical setup.



03 / Run and record

Clear the guidance, run the model, and link the result to the setup.

Light puts the instruction at the point of action. Cameras check that it happened.

Useful feedback with a fraction of the trials

Enact (YC)

Robotics lab

Before

With Prism

300–400 → ~30

Evaluations needed for useful engineering feedback

Before  300

With Prism  30

10× fewer evaluations

New Theory

Physical AI lab

Before

With Prism

40% → 3%

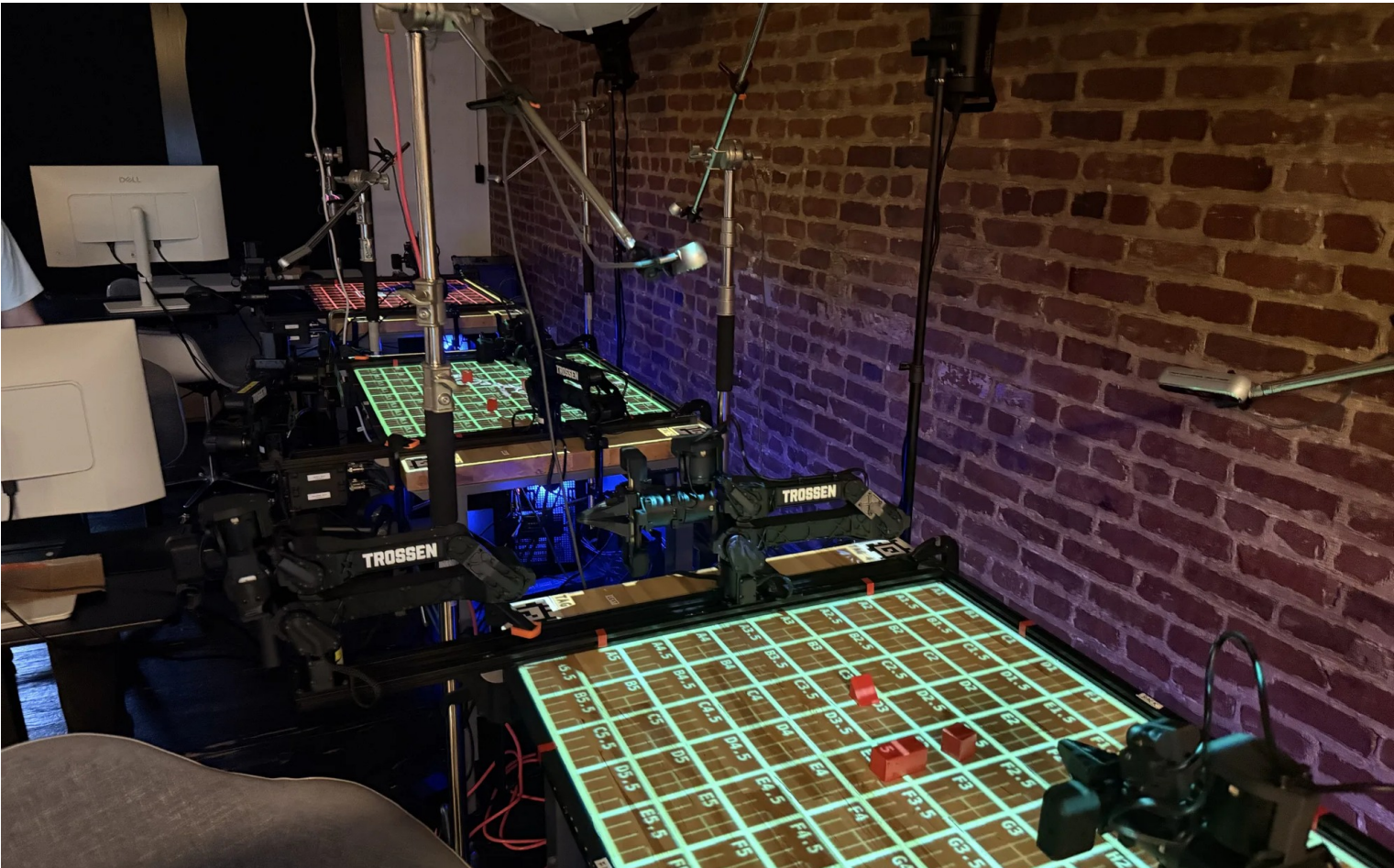
Evaluation variability over 30 days

Before  40%

With Prism  3%

92.5% relative reduction in variability

Ten paid units, already in the field



Prism installations in a working lab

10 paid installations across three customers. Customer use differs, as shown above.

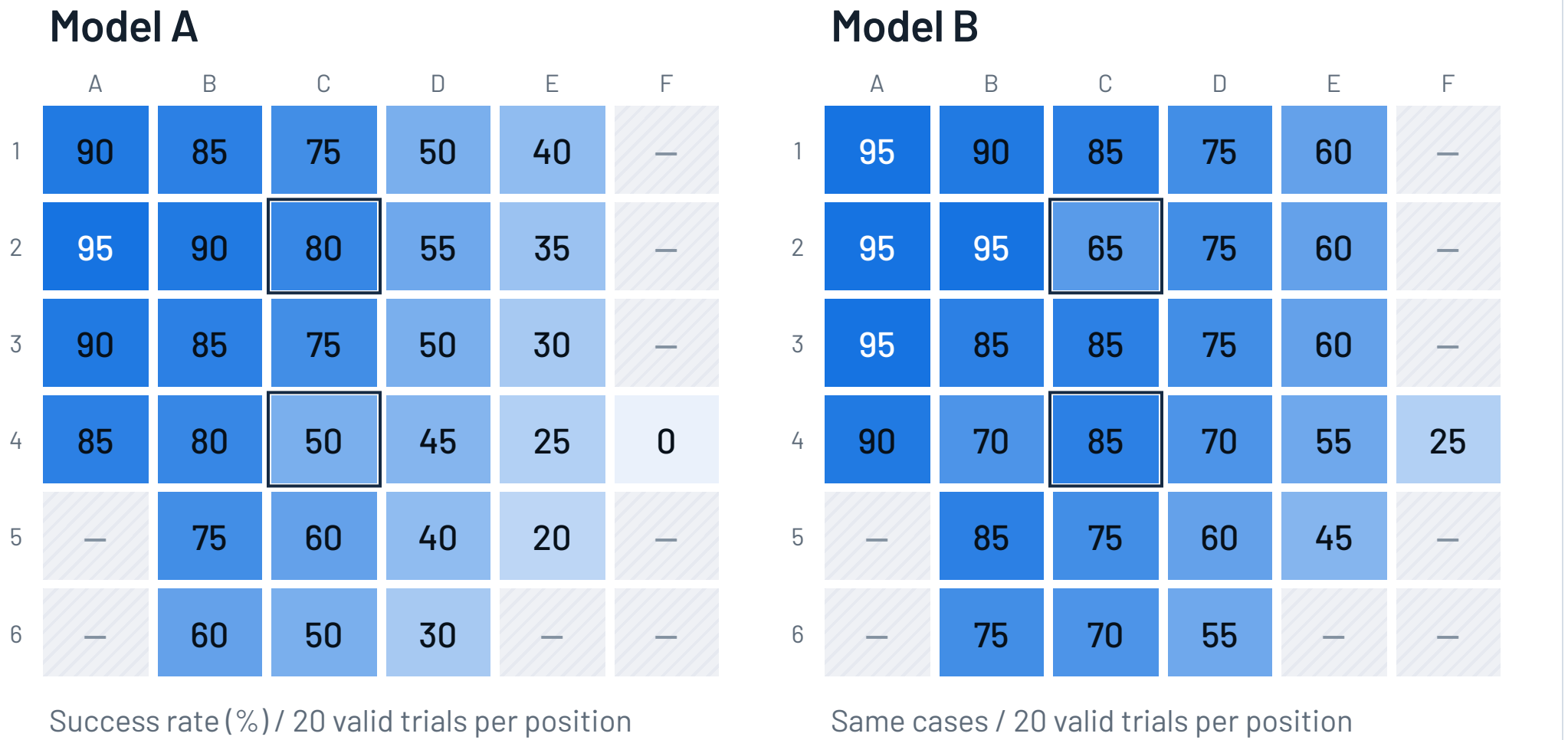
Enact Robotics evaluations	5
New Theory Robotics evaluations	4
Trossen Robotics Assessing Prism alongside its robots	1

Enact heard about us from another lab, tried one unit, and expanded to five.

We built and installed the units by hand. Repeatable production and setup are the next step.

Know where the next model improved

Compare the same physical cases. Catch regressions. Find conditions nobody tested.



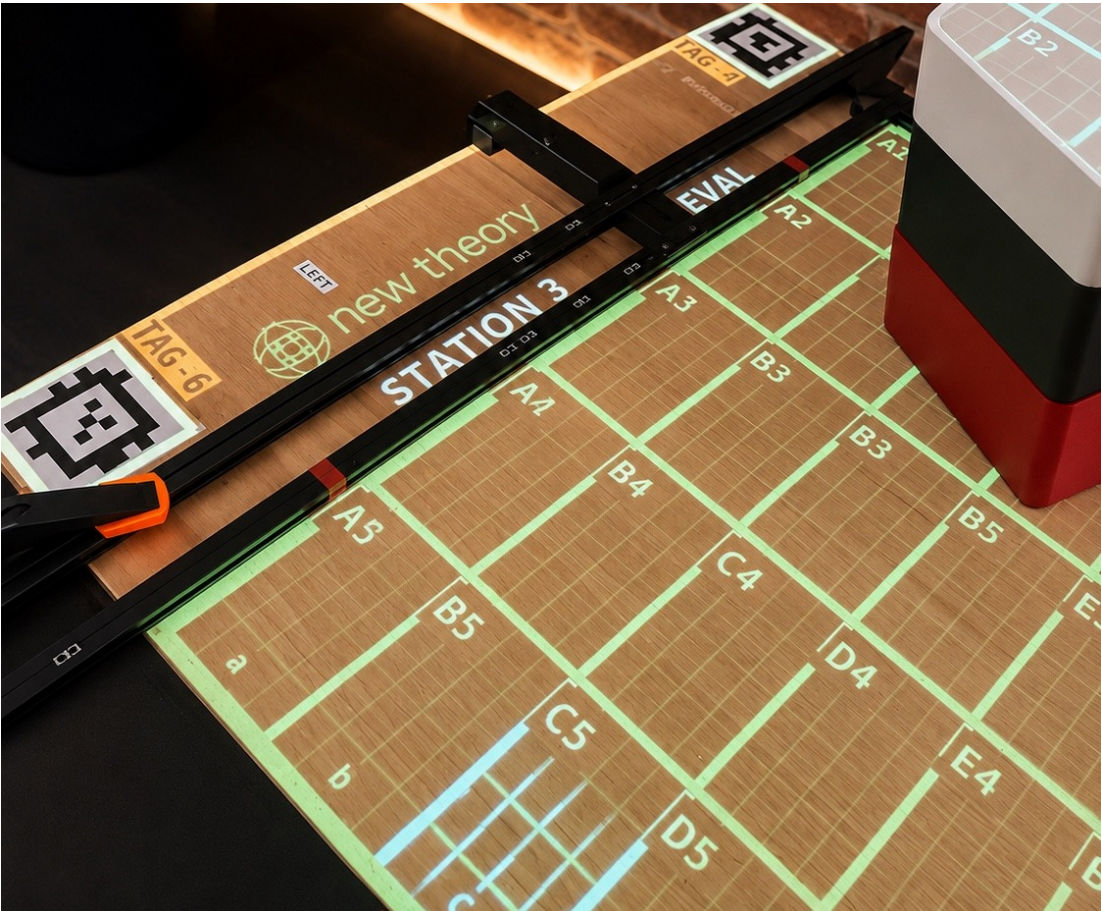
C4 / Improvement
50% to 85%
10/20 to 17/20 successful trials.

C2 / Regression
80% to 65%
16/20 to 13/20 successful trials.

F1 / Untested
No trials
A gap to investigate, not a failure.

Illustrative data, not customer results. Each tested cell has 20 valid trials per model. Numbers show success rate (%). Blank cells have no trials. Spatial maps are in expanding rollout.

A useful failure has a physical address



Real Prism projection. Case details at right are illustrative.

Example case / CUP-T0-TRAY-C4

Task + object	Place a matching cup fully inside the tray and release it.
Verified starting condition	x: 30 cm, y: 40 cm, orientation: 0° Position tolerance: 5 mm
Model + result	Model A: 10/20 successes Model B: 17/20 successes
Outcome + quality	Record the evaluator, images and validity flags. Keep rejected setups out of valid trial counts.

Preserve the conditions behind a failure, so the next engineer can recreate it.

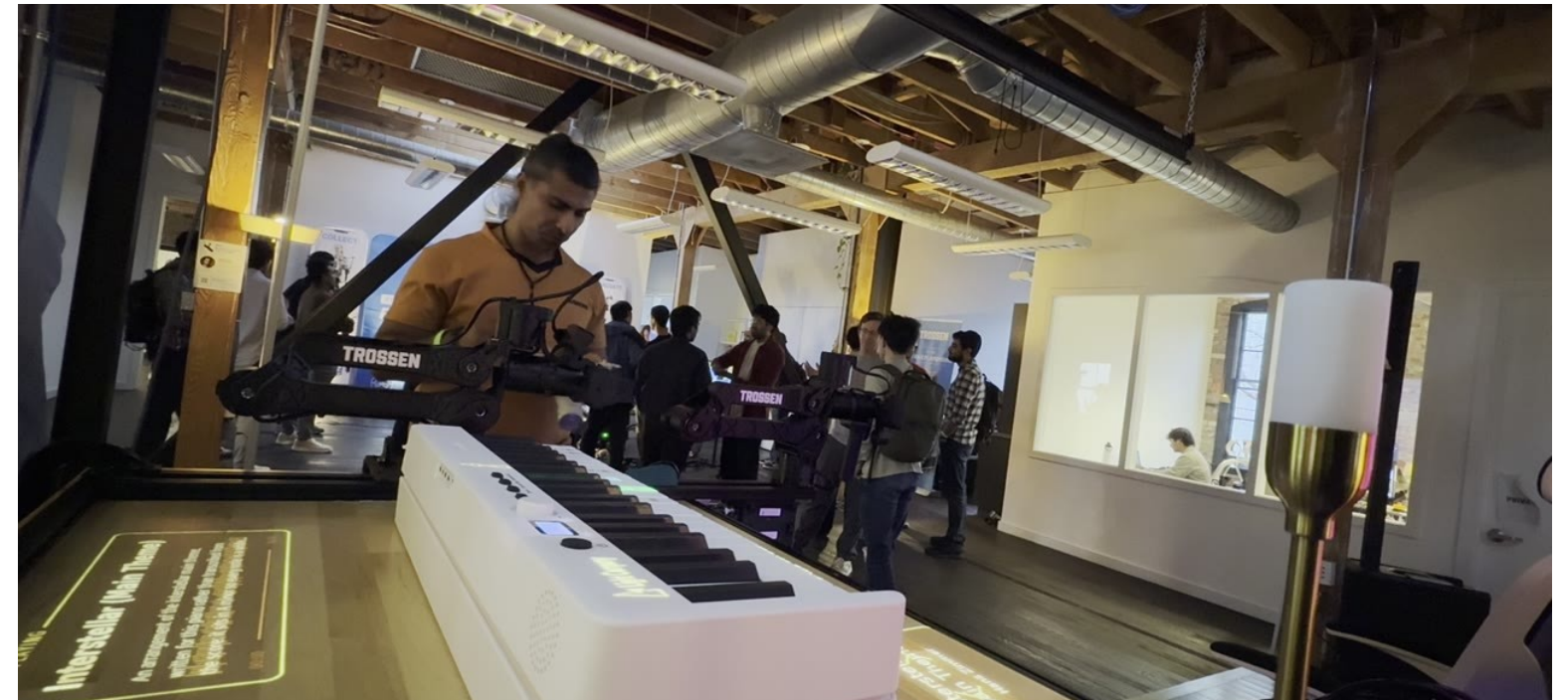
Example requirements, tolerances and results illustrate the data structure. A receiving station needs local calibration, compatible objects and a consistent success criterion.

The data opportunity grows beyond one lab

Customers pay us to install the infrastructure that produces this evidence.

The software helps a lab compare releases, find weak conditions, and choose its next tests. Contributed cases can reveal weaknesses another lab has not discovered.

Our ambition: the largest dataset of verified real-world robot evaluations.



Guidance and measurement inside the daily evaluation workflow

Customer's own history

Replay cases across models and stations.

Participating labs

Buy useful tests and benchmarks beyond their own workcells.

Broader ambition

Extend compatible case libraries across more robot hardware.

Shared benchmarks and data products are planned. Cross-customer reuse requires permission, compatibility checks and consistent scoring.

Pilot sales proved the unit economics.

\$1,500

pilot selling price per unit

\$300

Build + deploy

\$1,200

Contribution per unit

80%

pilot unit margin before
operating expenses



\$15,000

hardware revenue from 10 paid pilot units

Service charges add revenue beyond the hardware sale.

These were our pilot terms. Pricing will evolve across hardware, software usage and services.

Pilot unit margin = $(\$1,500 - \$300) \div \$1,500$. Unit margin excludes ongoing software, support, payroll and other operating expenses.

Software revenue can grow with usage and data

Our ambition

LMArena for robots.

Two revenue streams: evaluation usage and the benchmark/data products it makes possible.

Usage scenario

100 paid evaluations per station per day

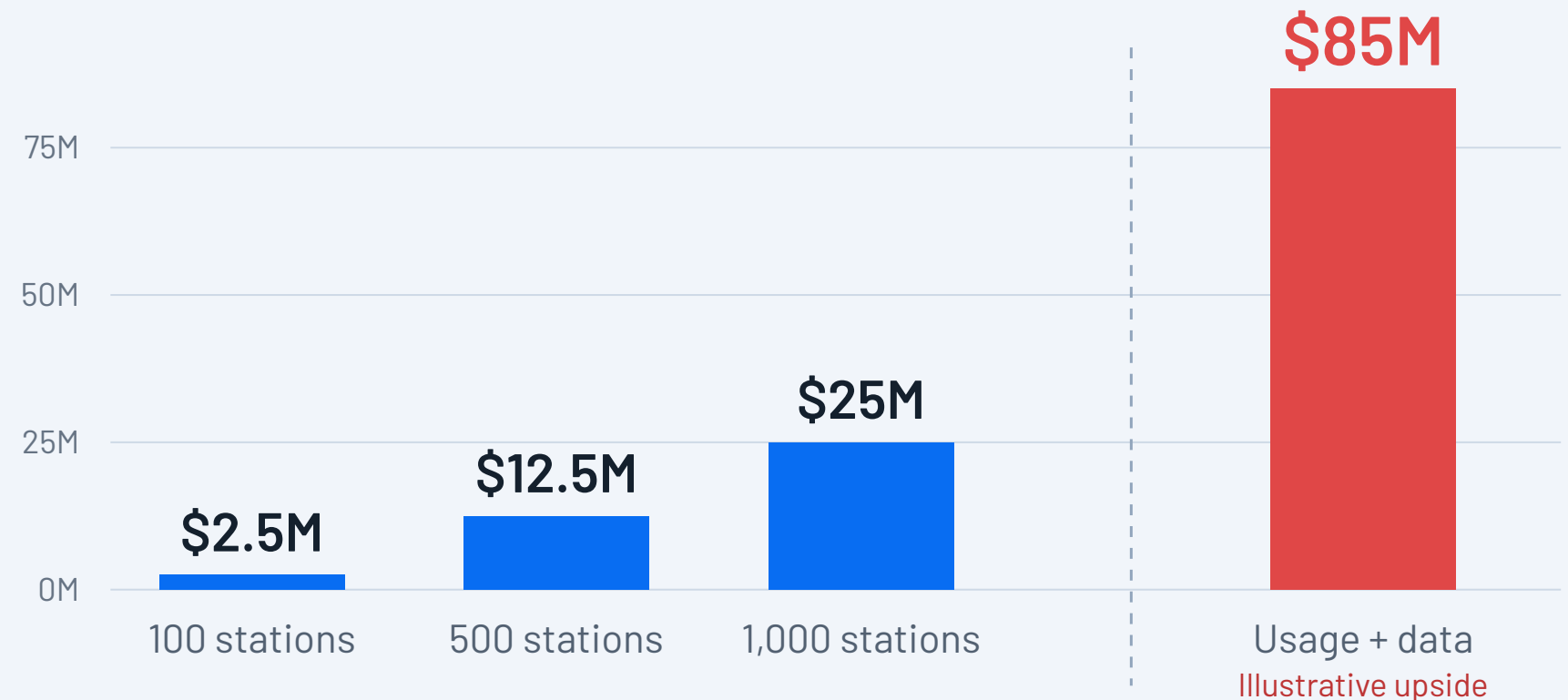
250 active days per year

\$1 per completed, verified evaluation

Software and usage pricing first.

Benchmark and data products next.

Illustrative annual revenue



1,000 stations × 100 evaluations × 250 days × \$1 = **\$25M**

\$85M = \$25M usage + \$60M benchmarks/data

The \$60M data-product revenue is an assumed scenario input.

Illustrative revenue scenarios, not current revenue or forecasts. Hardware sales excluded. Shared data products require participating-lab permission.

The next stage: scale deployment and prove paid software

100

paid stations

20

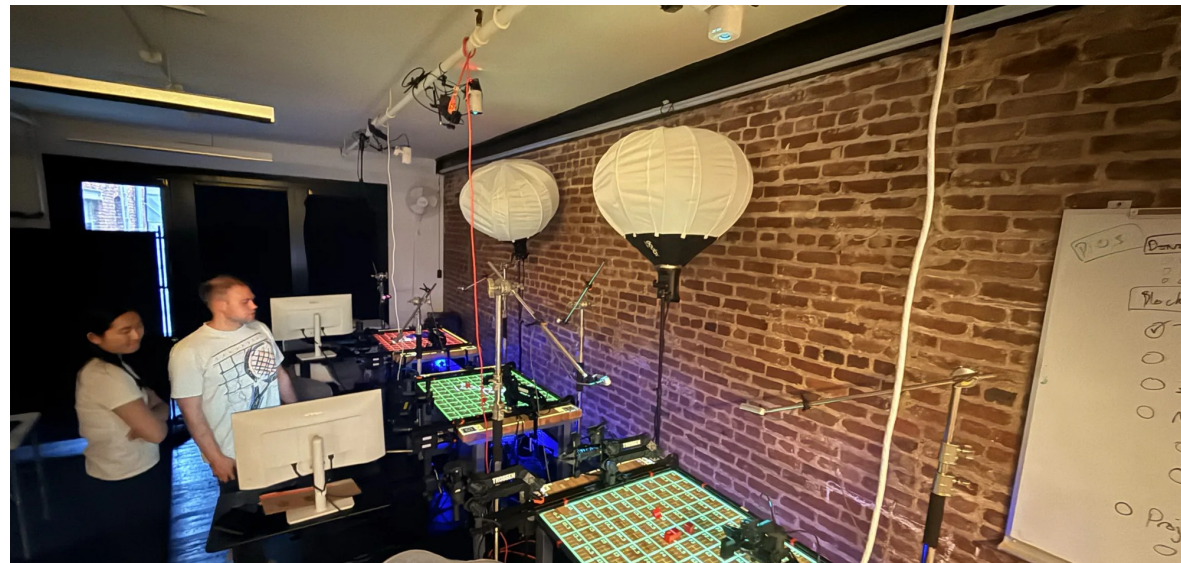
evaluation labs

10

labs paying for software

50,000

valid evaluation episodes



Building beside the operators who need it

Make installation repeatable

Standardize the unit, secure dependable parts, and let customers install and calibrate without us onsite.

Prove the data is useful beyond one lab

Launch a paid failure-case pilot. Measure whether a contributed case exposes a previously untested weakness in another compatible lab.

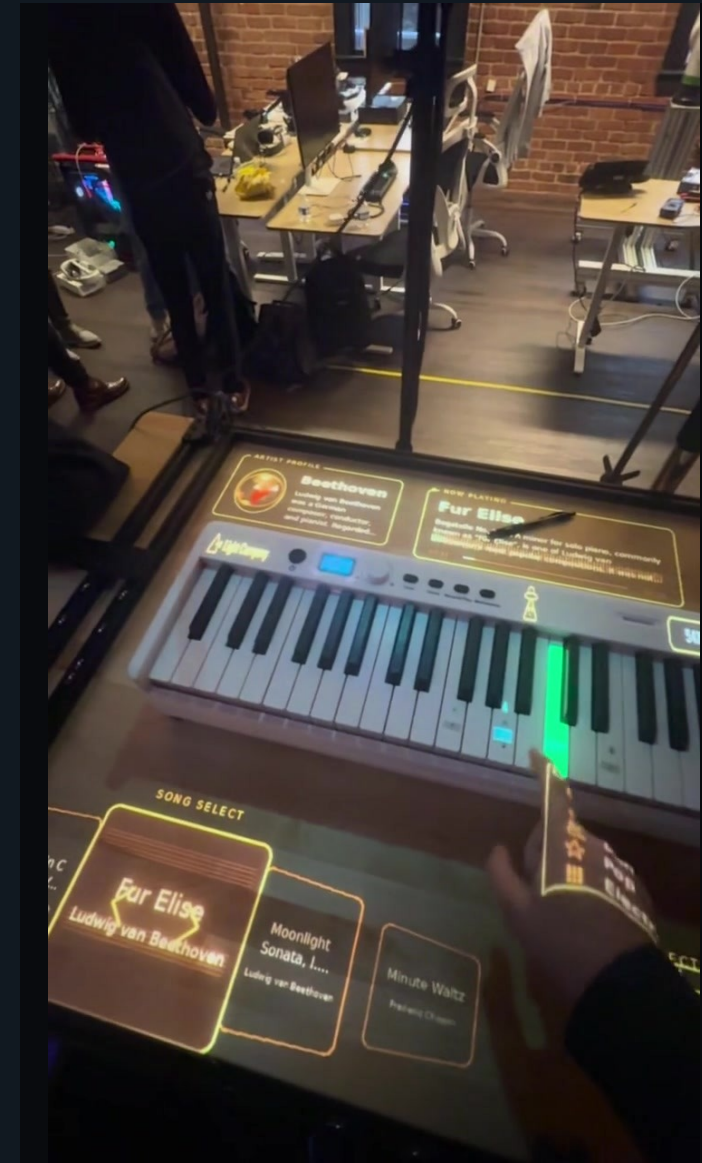
The original vision: AR without glasses.

Intelligent light that understands the workspace, guides the next action, and checks the result.

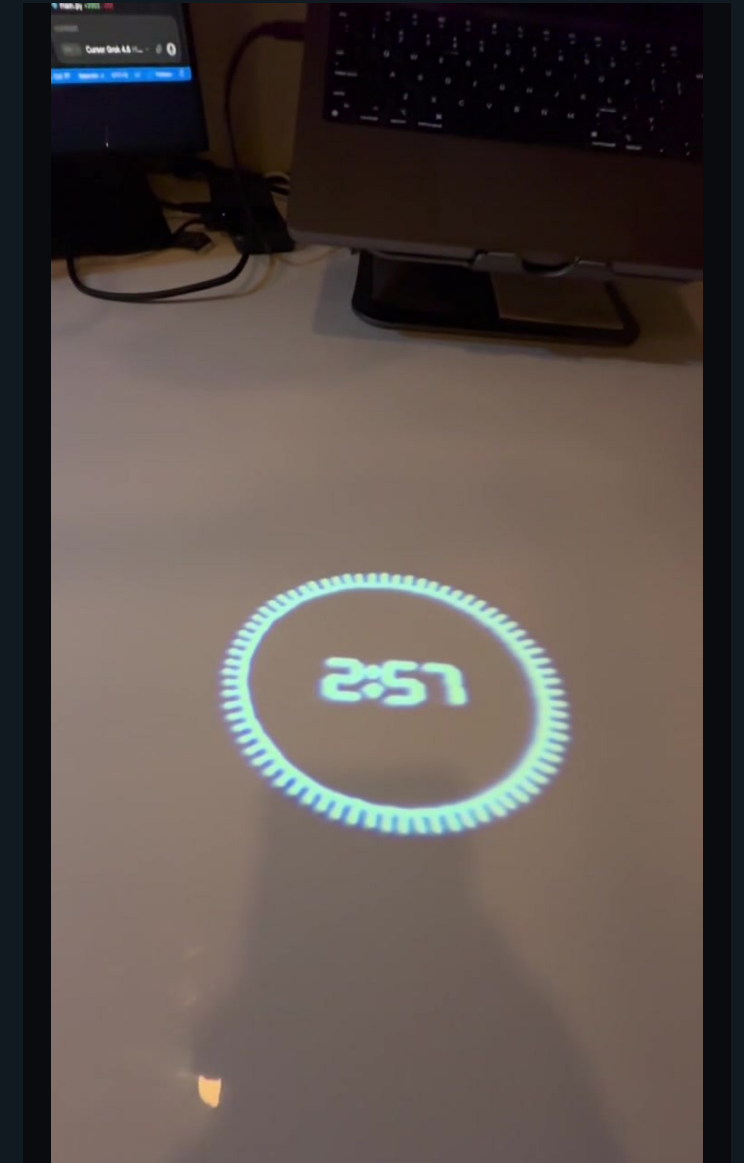
Robotics evaluation is our first commercial use. The same interaction can help people learn, assemble, repair and work in the physical world.

[Explore robotics ↗](#)

lightcompany.ai/demo ↗ hi@jonny.sh



Guidance on real piano keys



Information on the work surface