



Breacher.ai

DEFEND WITH KNOWLEDGE

FINDINGS FROM AUTHORIZED ENGAGEMENTS

Click rate does not predict outcomes.

Four findings from orchestrated social engineering engagements, scored by outcome instead of behaviour, and what that changes about where the next control dollar goes.

What the outcome data says



01

Click rate does not predict outcomes

The function that acted most often produced no consequence at all. The functions that produced real consequences behaved well on every conventional metric.

02

Awareness training reduces risk

It moves the number it was built to move. Trained populations carry lower SPREAD, fewer people behave the wrong way, and the cheaper half of the problem shrinks.

03

The training platform does not predict the outcome

Which awareness platform an organization runs shows no relationship to DEPTH. Consequences landed, and were contained, across platforms alike. The control path decides, not the curriculum.

04

Technology was the riskiest sector by outcome

The best-behaving people in the book, and the only sector where consequences actually landed. Rank sectors by failure rate and Technology looks safe.

THE SCALE

One score, two numbers

A single figure hides which of the two is driving it. The split is the finding.

DEPTH · 0-70

Did anything actually get through?

Set by the furthest point any single person reached. Never divided by headcount, one person moving money sets it for the whole organization.

SPREAD · 0-30

How much of the population behaved that way?

Weighted incidence across everyone tested. Behaves like a conventional rate, capped so it can never outweigh a real outcome.

BANDS

Low	0-39 Nothing reached an action
-----	-----------------------------------

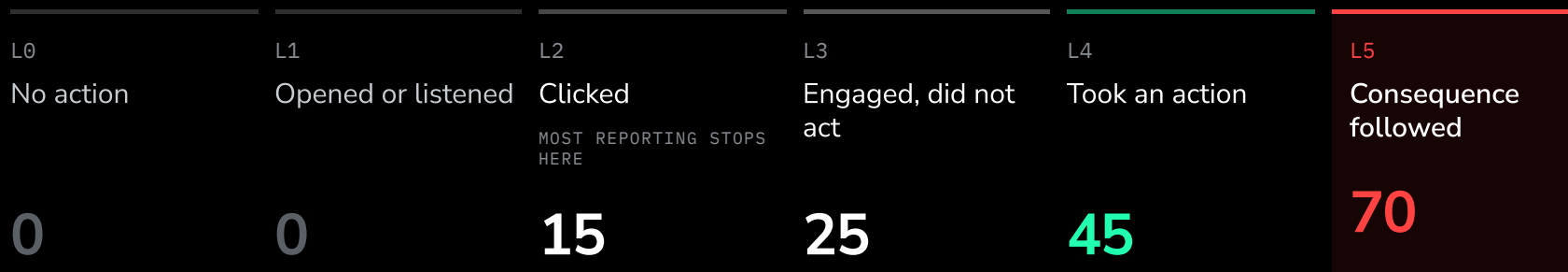
Medium	40-59 Someone acted
--------	------------------------

High	60-79 Widespread action, or a consequence that landed narrowly
------	---

Critical	80-100 A consequence landed and the behaviour was not isolated
----------	---

THE LADDER

What sets DEPTH



An **action** is credentials submitted, a payment initiated, access granted or data sent. **L5** means nothing downstream caught it. Your training stops at the click. The attack does not.

FINDING 01

Click rate does not predict outcomes

The two numbers move independently. Ranking on the visible one inverts the ranking on the one that costs money.

WORST ON BEHAVIOUR · BEST ON OUTCOME

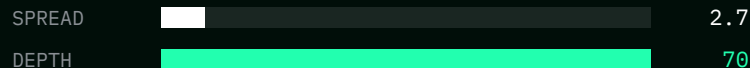
IT



The widest action of any function, and not one consequence. Every instance was contained downstream. On a click metric, your worst department. On outcome, your best evidence that controls are working.

BEST ON BEHAVIOUR · WORST ON OUTCOME

Technology



Almost nobody behaved badly, and consequences still landed: a credential that worked, a payment that moved, a synthetic candidate that cleared a hiring screen. The exposure is not in the people. It is in what sits behind them.

Same engagements, opposite rankings. A click-rate report would put IT on the remediation list and sign Technology off as a good result.

Awareness training reduces risk

It moves the number it was built to move. That is a real reduction, on one of the two axes.

WHAT TRAINING MOVES

SPREAD

falls

Fewer people click, fewer engage, fewer act. Trained populations sit visibly lower on incidence, and the weakest-behaving sectors are the ones where the training and process work has not been done.

DEPTH

unchanged

DEPTH is set by whether a control caught the action, not by whether the person recognised the lure. Training cannot close a path with no control on it.

THE EVIDENCE IN THIS BOOK

24.8 Manufacturing SPREAD. The weakest behaviour we measure, and the cheaper problem to fix.

1.7 Legal SPREAD. The strongest sector in the book. Fifteen times less incidence than Manufacturing.

45 Both carry the same DEPTH. Behaviour separated them by 23 points of SPREAD and zero of outcome.

Keep training. Stop scoring it as the whole programme. High SPREAD, low DEPTH is a training problem, the cheaper one. High DEPTH, low SPREAD is a control problem, and the one that produces the loss.

The platform does not predict the outcome either

WHAT PLATFORMS COMPETE ON

Content, cadence, completion

Every differentiator in the category sits on the SPREAD axis: how many people behave the wrong way in an email test. Real differences exist there. None of them reach DEPTH.

WHAT ACTUALLY SET DEPTH

A missing control on one path

Replayable credentials. A payment approval seniority could waive. A hiring pipeline with no identity check. Each consequence traced to a control, never to a knowledge gap.

WHAT THIS IMPLIES

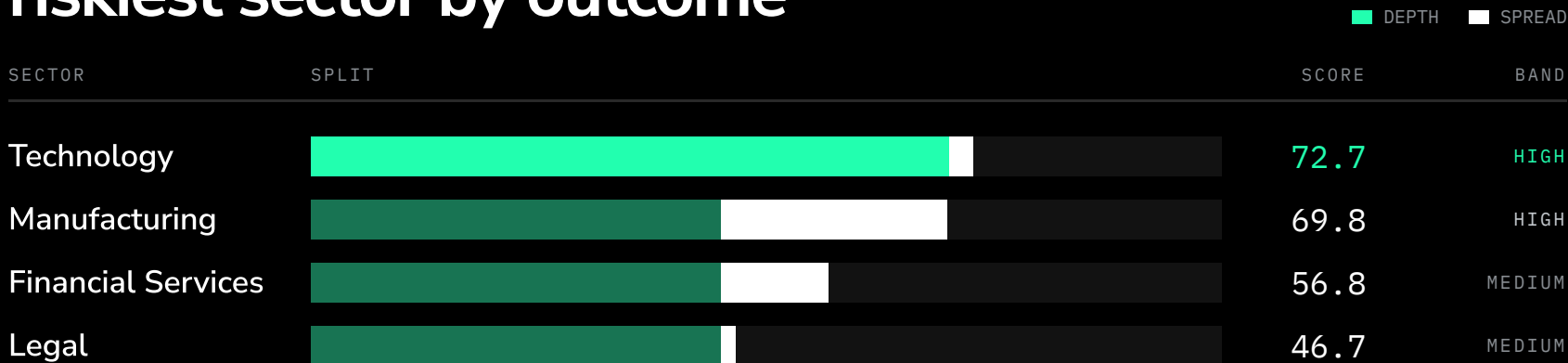
Switching vendors is not a control change

A platform migration moves SPREAD at best. If your DEPTH is 70, it will still be 70 the quarter after the migration, because nothing in the loss path changed.

Every consequence in this book involved synthetic voice or video. None was a plain email, which is the only channel most platforms can test.

FINDING 04

Technology was the riskiest sector by outcome



The top two are opposite problems. Technology sits highest on the strongest behaviour we measure; Manufacturing sits second on the weakest. Rank these sectors by how often people fail and you get the order backwards on what it cost.

Every action that had a consequence

	FUNCTION · CHANNEL	WHAT HAPPENED
01	Human Resources SYNTHETIC VIDEO INTERVIEW	A synthetic candidate progressed through a hiring screen and reached the point of system access. No security control existed anywhere in the path.
02	Finance SYNTHETIC EXECUTIVE VIDEO	A payment instruction delivered by synthetic executive video resulted in a wire transfer being initiated.
03	All-staff REPLICA SIGN-IN PAGE	Corporate credentials were submitted to a replica sign-in page.

Humans cannot spot them

Deepfake detection assessment · 7 samples per participant

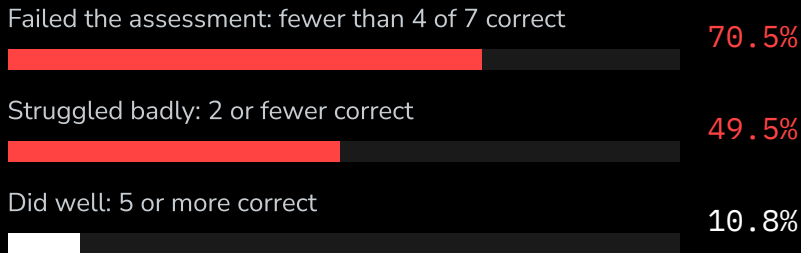
WORSE THAN A COIN FLIP

38%

average accuracy · 2.66 of 7 correct

Random guessing scores 50%. Below that line, people are not merely failing to detect a deepfake. They are being actively deceived by it.

PERFORMANCE BREAKDOWN



Roughly one in two got two or fewer right. About one in ten can detect reliably, and there is no way to know in advance which of your people that is.

Training people to spot deepfakes is a flawed control. Detection is not a skill you can reliably teach at scale. Put the control in the process (verification, callbacks, authentication) where it does not depend on someone recognising a face or a voice.

FUNCTION

Where it actually sits

CRITICAL

Human Resources

88.4 70 + 18.4

A synthetic candidate cleared a hiring screen with no security control anywhere in the path. Recruitment is not modelled as an attack surface.

Act on: identity verification in hiring, and a control where a candidate first gains access.

CRITICAL

All-staff

83.0 70 + 13.0

Good behaviour, real outcome. Corporate credentials were submitted to a replica sign-in page and they worked.

Act on: phishing-resistant authentication: a credential that cannot be replayed.

CRITICAL

Finance

80.4 70 + 10.4

A payment instruction delivered by synthetic executive video resulted in a wire transfer being initiated. Small, high-authority group, so incidence will always understate it.

Act on: callback on banking changes, and dual approval seniority cannot waive.

HIGH

IT

66.5 45 + 21.5

The widest action of any function and no consequence at all. People acted, and every instance was contained.

Act on: capture why each action failed to escalate, the most valuable unrecorded data you hold.

Three of four functions are Critical for the same reason. HR, Finance and the all-staff population each carry a DEPTH of 70 because each produced a real outcome. None has a broad behavioural problem. IT, with the widest action of anyone, is the only function not Critical.



Breacher.ai

DEFEND WITH KNOWLEDGE

Get your own score

A scoping call takes thirty minutes: which functions to test first, which channels matter for your threat model, and what a baseline would surface that your current reporting cannot.

BOOK	cal.com/breacher.ai
THE MODEL	breacher.ai/social-engineering-risk-index
CONTACT	support@breacher.ai · +1 407 900 0799

Orchestrated

Email, voice, messaging, video and live conversation run as one coordinated campaign, the way a real operator works, not five disconnected tests.

Fully managed

Run externally. Nothing installed, no integration into your stack.

OSSES™ Risk Score, per function

Every score in this deck is an OSSES™ Risk Score: DEPTH and SPREAD by department, so you can see which of the two problems you actually have.

Consent-governed

Executive likeness and voice simulation runs only under documented, signed authorization.