

10 questions to ask before you trust someone else's intel

Threat intel is only worth paying for if it changes what's detecting in your environment. Ask about it while you're still evaluating, when the answers are easiest to get. The best answers have a date, name, or number in them (or better yet, all three).

01 Where does your intel come from, and what's the mix?

A good answer: Their own incidents first, outside feeds second, and they'll give you the ratio.

Red flag: A list of feeds they subscribe to. You can buy those without them.

02 When did you last turn something new into a detection, and how fast?

A good answer: A specific story with a timeline in hours or days, and the team who did it.

Red flag: A description of "our process," or a timeline measured in release cycles.

03 Which threat groups have you named yourselves?

A good answer: Ones they identified first, with published research behind the name.

Red flag: Every name they use came from someone else's research.

04 Do the people writing your intel also work incidents?

A good answer: Yes, and close enough that a Tuesday finding changes coverage by Thursday.

Red flag: A research team that makes marketing content and never touches your environment.

05 Can I see the detection logic?

A good answer: Yes, in the platform, whenever you want. What fired, what didn't, and why.

Red flag: It's intellectual property you can't inspect. You're buying an unauditable outcome.

06 How fast does what you learn from one customer reach the rest?

A good answer: Automatically, and fast. Coverage built in one environment ships to everyone.

Red flag: It stays with the account team that found it, or arrives as a PDF next month.

07 What do you do with intel besides send me a report?

A good answer: New detections, hunt hypotheses, and fixes specific to your environment.

Red flag: The intel program is a content program. Reports arrive, nothing changes.

08 A critical CVE drops today. When do I hear from you?

A good answer: A sequence with timing: who assesses it, how they check you, how you find out.

Red flag: You hear it from the news, then from their newsletter.

09 How do you hunt for what your detections miss?

A good answer: Hypothesis-based hunts with written findings. Ask what the last one found.

Red flag: "Hunting" is a search bar in your platform, or a report you can't trace back to a hypothesis.

10 Do you tell me when something matters, or only when there's an action for me?

A good answer: They flag what they're seeing as they see it, whether or not you have to act.

Red flag: Notification only after they've confirmed something is actionable.

A little intel on Expel Threat Intel

Expel Threat Intel is built on what our security operations team sees. Analysts, detection engineers, and researchers work the same incidents, so a finding in one customer's environment becomes coverage for everyone.



Expel's intel comes from our own incidents first. 2,450+ proprietary detections, 160+ coverage options.



All detection logic is visible to customers in Expel Workbench™.



Expel names and tracks its own threat groups, with published research behind them.



Expel offers hypothesis-based threat hunting, with written findings you receive.

Come hear how Expel answers these questions.

Join our next Democast, a 30-minute group demo and live Q&A.

expel.com/democast