

Your Data in the Age of AI Agents

How AI browser agents have changed the shape of data-exfiltration risk for buying groups — and how the eBiz platform now detects suspicious activity, machine-driven bulk extraction and out-of-character behaviour while it is still happening.

Your platform holds some of your organisation's most commercially sensitive information — negotiated deal terms, rebate structures, supplier agreements and price lists. That information is valuable: to competitors, to other suppliers, and to employees preparing to leave.

The eBiz platform is secure, regularly penetration tested, and has never been hacked or compromised. But the uncomfortable truth about data exfiltration is that it almost never looks like a “hack”. It looks like a **legitimate login doing something unusual**: a user who normally views three or four deals a day suddenly opening two hundred, or downloading every price list in the system over a weekend. Because the access itself is authorised, traditional security tools — firewalls, password policies, intrusion detection — see nothing wrong. By the time anyone notices, the data is already gone.

AI browser agents have made that risk sharper. Tools such as Anthropic’s Claude for Chrome, OpenAI’s browser agents and their open-source equivalents let any user hand control of their logged-in browser session to an AI. An agent can quietly walk an entire agreement library in hours, extracting deal and rebate data record by record and compiling it into a single spreadsheet. To the system, it just looks like a very busy user.

Every hour

the platform reviews the last 24 hours of activity for every user, against their own 30-day baseline

4× norm

the default multiplier of a user’s personal daily activity that triggers a suspicious-activity alert

2+ signals

independent machine-behaviour patterns that must coincide before an account is flagged as an AI agent

Figure 1 — Why exfiltration doesn’t look like a hack

Every step uses valid credentials and ordinary pages. Perimeter security sees a normal session; only behaviour gives it away.



What perimeter security sees at each step:

“a valid user, logged in, clicking through ordinary pages”

What behaviour shows: a pattern no human produces.

This paper explains the risk in plain terms, then sets out the layered defences now built into the platform: behavioural monitoring, AI-agent detection, immediate evidence-rich alerting, and customisable usage terms that put your protection on a contractual footing.

PART ONE

A very busy “user”: what AI agents change

Platforms like eBiz have always used **“friction as security”** alongside access controls: there is deliberately no button that downloads every agreement and rebate in one go. A person who wanted the whole library had to click through it record by record — slow, tedious, and rarely worth it.

AI browser agents remove that friction without breaking any rule the perimeter can see. Given a one-line instruction in plain English, an agent will navigate the site, apply filters, open each record, extract the values that matter — including rebate levels — and assemble the results into a tidy spreadsheet. It does the tedious work at machine pace and never gets bored.

Two things make this different from the scrapers of the past. First, the agent operates **inside a genuine, authorised login** — there is no stolen password and no breached firewall to detect. Second, it can behave convincingly like a person when the task is small, which means **no single tool can be the whole answer**. Detection has to be layered.

CONTROLLED TEST – LIVE SITE, JULY 2026

Working with one of our customers, an administrator asked a consumer AI browser agent, in one sentence, to *“list all calculated rebates”* for a single product category — using nothing more than a standard user login.

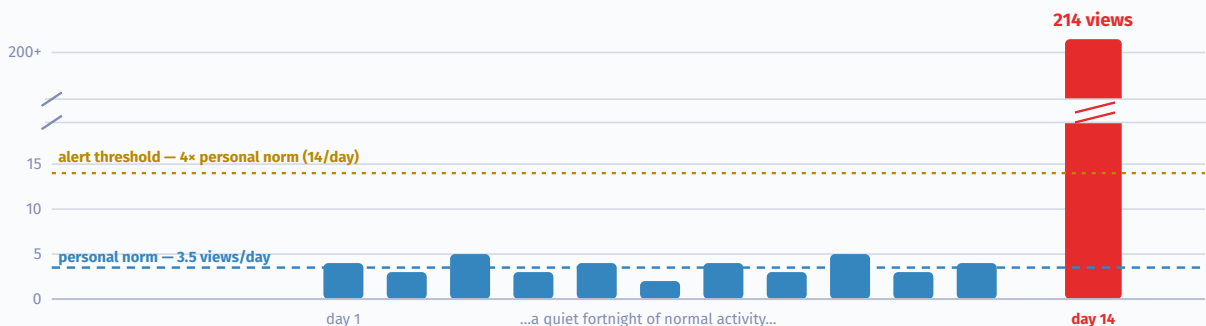
The agent navigated the menus, found the right screens, applied the right filters and returned an accurate, multi-page breakdown of rebate levels **within minutes**. No security control was bypassed, because none was triggered: to the platform, the session was an ordinary authorised user.

The same test is why the defences described in this paper exist. The gap is real, and it is now being watched.

The insider dimension matters as much as the AI one. The classic exfiltration event is not an outsider at all — it is a legitimate account, often in the weeks before a resignation, quietly collecting deals, price lists and agreements. The detectors below treat both the same way: **it is the behaviour that is suspicious, not the tool**. A human who bulk-downloads 32 agreements back-to-back gets flagged just as an agent does.

Figure 2 — What “out of character” looks like

Daily deal views for one account over 14 days. The baseline is the user’s own 30-day history — not a global average. Alert threshold shown at 4× the personal norm (default).



Illustrative, using the format of a real alert: “214 deal views in 24h vs a typical 3.5/day” — roughly 61× this user’s personal norm. A genuinely heavy user being their usual heavy self never crosses their own 4× line; a quiet account suddenly bulk-reading the library does.

PART TWO

Two detectors, watching different things

The platform's Suspicious Activity module runs two independent detectors around the clock against the full audit trail. One watches **volume against personal history**; the other watches for **patterns humans don't produce**. Both are live now on the platform.

1 Suspicious activity detection

Flags accounts whose deal viewing or document downloading is out of character *for them personally*.

- **Every hour**, the last 24 hours of activity is reviewed for every user.
- Users above a minimum volume are compared to their **own 30-day baseline**.
- An alert fires only at **several times the personal norm** (4× by default).
- **New users are never flagged** — an account needs 14+ days of history and 5+ active days before it can be judged “out of character”.
- No hourly re-alert spam — but continued activity triggers a **“continued suspicious activity” escalation**.

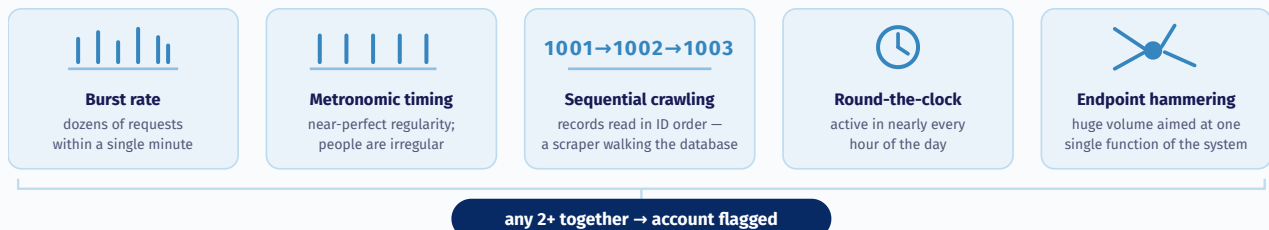
2 Agent detection

Flags accounts behaving in ways humans don't — across *all* platform activity, not just deals and documents.

- Watches five independent, tell-tale **machine patterns** (Figure 3 below).
- **No single quirk triggers an alert** — a flag is raised only when two or more independent signals coincide, keeping false positives low.
- Deliberately applies to **brand-new accounts too**: a day-old login behaving like a script is more suspicious, not less.
- Catches automation regardless of the tool — consumer AI agents, scripts, or scrapers.

Figure 3 — The five machine signals

Independent behavioural fingerprints of automation. Any two or more together raise a flag.



What lands in your inbox when something is flagged

When either detector fires, a security alert email goes immediately to your nominated recipients — while the activity is still ongoing, not weeks later:

Who

Name, company, job title, email address and account ID of the user concerned.

Why — in plain English

The specific evidence, e.g. “214 deal views in 24h vs a typical 3.5/day” or “sequential ID crawl: 45 consecutive ascending IDs”.

The numbers

Current activity side-by-side with the user's own historical baseline, so the scale of the anomaly is obvious.

The full 7-day activity report

Attached as a PDF with a live link: every deal viewed and document downloaded (named, timestamped), daily and hour-of-day charts with out-of-hours activity highlighted, connecting IP addresses, and error counts.

Every alert is written permanently to the platform's audit feed, and the same activity report can be run **on demand against any user at any time** by an administrator — you don't need to wait for an alert to investigate. All thresholds are configurable per site under **Settings > General > Suspicious Activity Alerts**.

PART THREE

Defence in depth — and terms with teeth

Customisable usage terms LIVE

Detection tells you something happened; **terms give you standing to act on it**. Standard website terms are a tick-box exercise and won't protect you against the behaviour described in this paper. The platform therefore now supports usage terms that **you define** — specific to your business and buying group — displayed on the login screen and accepted before anyone gains access. Combined with the audit trail and activity reports, that gives you real evidence and a far stronger basis for recourse if someone acts outside them.

Our position on AI UNCHANGED

None of this is anti-AI. Used well, AI assistants genuinely help your members and your business, and many boards are actively encouraging adoption. The point of these defences is narrower: **your commercially sensitive data should leave the platform only on your terms** — whether the thing moving it is a person, a script or an AI agent. It is the behaviour that gets flagged, never the tool.

Network-edge bot management IN EVALUATION

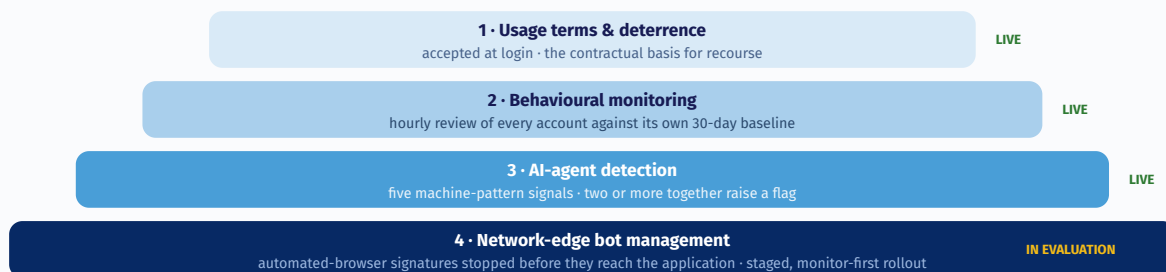
Detection inside the platform is being complemented by a further layer at the network edge, currently in evaluation: managed bot-control tooling (such as AWS WAF Bot Control's targeted protections) that recognises the signatures of automated browsers **before a request even reaches the application**. Rollout is deliberately staged — running in monitor-only mode first, so legitimate traffic is never disrupted — with application-layer session-token validation planned as a second line behind it.

What we recommend you do now

- 1 Nominate your alert recipients** — the people who should receive security alerts the moment a detector fires.
- 2 Tune the thresholds with us** — we'll walk through the defaults (4× norm, minimum volumes, signal sensitivity) in your regular call.
- 3 Draft your usage terms** — we'll help you write terms that fit your group and put automated bulk extraction explicitly out of bounds.
- 4 Brief your members** — the quiet deterrent: bulk extraction is now detected while it happens, evidenced, and traceable to a named account.

Figure 4 — The layered defence model

No single layer is the whole answer; an agent that slips one layer meets the next.



These features are live on your platform now.

To tune your thresholds, nominate alert recipients or draft your usage terms, raise it in our next regular call — or contact us directly.

eBiz
Hove Business Centre, Fonthill Road, Hove BN3 6HA
ebiz.uk