

Security Evaluation Test Report

Enterprise Endpoint Security



ONLINE REPORT

SE LABS tested a variety of anti-malware (aka ‘anti-virus’; aka ‘endpoint security’) products from a range of well-known vendors in an effort to judge which were the most effective.

Each product was exposed to the same threats, which were a mixture of targeted attacks using well-established techniques and public email and web-based threats that were found to be live on the internet at the time of the test.

The results indicate how effectively the products were at detecting and/or protecting against those threats in real time.

Contents

Introduction	04
Executive Summary	05
Security Evaluation EPS Protection Enterprise Awards	06
Threat Responses	07
1. Protection and Legitimate Handling Accuracy	08
1.1 Protection Details	08
1.2 Attack Types	08
1.3 Total Accuracy Ratings	09
1.4 Protection Accuracy	09
1.5 Protection Scores	10
1.6 Legitimate Accuracy Ratings	10
2. Conclusion	11
Appendices	12
Appendix A: Protection Ratings	12
Appendix B: Legitimate Interaction Ratings	13
Appendix C: Terms Used	15
Appendix D: FAQs	15
Appendix E: Product Versions	16

Document version 1.0 Written 27th July 2026



Management

Chief Executive Officer Simon Edwards

Chief Operations Officer Marc Briggs

Chief Human Resources Officer Magdalena Jurenko

Technical Director Thomas Bean

Testing Team

Nikki Albesa

Daniel Botez

Solandra Brewster

Billy Coyne

Jarred Earlington

Gia Gorbald

Anila Johnny

Cameron Love

Jeremiah Morgan

Julian Owusu-Abrokwa

Joseph Pike

Enejda Torba

Marketing

Sara Claridge

Isobel Edwards

Publication

Rahat Hussain

Colin Mackleworth

IT Support

Danny King-Smith

Chris Short

Website selabs.uk

Email info@SElabs.uk

LinkedIn www.linkedin.com/company/se-labs/

Blog <https://selabs.uk/blog/>

Post SE Labs Ltd,

55A High Street,

Wimbledon,

SW19 5BA, UK

SE Labs is ISO/IEC 27001 : 2022 certified and
BS EN ISO 9001 : 2015 certified for The Provision
of IT Security Product Testing.



© 2026 SE Labs Ltd

Introduction



CEO
Simon Edwards

If you spot a detail in this report that you don't understand, or would like to discuss, please contact us. SE Labs uses current threat intelligence to make our tests as realistic as possible. To learn more about how we test, how we define 'threat intelligence' and how we use it to improve our tests please visit our [website](#) and follow us on [LinkedIn](#).

Why Endpoint Security Misses Targeted Attacks

Effective protection depends on recognising more than malware

Recent targeted campaigns show that attackers still don't need novel malware or an undisclosed vulnerability to infiltrate an organisation.

Increasingly, they combine familiar tools, credible social engineering and legitimate system functions into an attack chain where each action may appear relatively harmless. The sophistication lies not in the malware used, but in the sequence of events.

An attack might begin with a convincing email or phone call, followed by a request to join a screen-sharing session. The victim may then be persuaded to install a legitimate remote-management tool, open a trusted application or run commands supplied by someone claiming to provide technical support. Each step can appear perfectly reasonable in isolation. Only when viewed together does the malicious objective become clear.

This creates a difficult problem for endpoint protection. Blocking every administrative tool, script or remote connection would make normal work impossible. But allowing them without sufficient scrutiny gives attackers opportunities to misuse trusted software and standard system functions. The challenge is distinguishing ordinary activity from behaviour that forms part of a developing attack.

This quarter's testing again showed that most products were highly capable when faced with widespread public malware. But some still struggled when attacks progressed through a targeted sequence of actions.

Timing matters. Detecting one suspicious action is useful, but detection alone may not prevent compromise. A warning that arrives after credentials have been stolen, persistence established or sensitive files collected is less valuable than an intervention that stops the sequence early. Equally, disrupting one stage may not be enough if the attacker can continue through another route.

Targeted attacks therefore provide a demanding test of endpoint protection. They require products to interpret context, connect related events and respond before a chain of individually plausible actions reaches its intended outcome.

Our endpoint security reports test more than a product's familiarity with a malicious file. We assess whether it can identify a developing attack, intervene decisively and prevent a successful first step from becoming a complete system compromise.

Executive Summary

Product Names

It is good practice to stay up to date with the latest version of your chosen endpoint security product. We made best efforts to ensure that each product tested was the very latest version running with the most recent updates to give the best possible outcome.

For specific build numbers, see **Appendix E: Product Versions** on page 16.

Products Tested	Protection Accuracy Rating (%)	Legitimate Accuracy Rating (%)	Total Accuracy Rating (%)
Broadcom Symantec Endpoint Security	100%	100%	100%
Kaspersky Endpoint Security	100%	100%	100%
CrowdStrike Falcon	100%	100%	100%
Trellix Endpoint Security	99%	100%	100%
Sophos Endpoint Security	98%	100%	99%
Microsoft Defender Antivirus (enterprise)	96%	100%	99%

● Products highlighted in green were the most accurate, scoring 85 per cent or more for Total Accuracy. Those in yellow scored less than 85 but 75 or more. Products shown in red scored less than 75 per cent.

For exact percentages, see **1.3 Total Accuracy Ratings** on page 9.

● The endpoints were effective at handling general threats from cyber criminals...

All products were very capable of handling email- and web-based threats such as those used by criminals to attack Windows PCs, tricking users into running scripts that download and run malicious files.

● ...but targeted attacks caused problems for some of the products.

Most of the products proved highly effective against the targeted attacks used in this test, while a couple missed one or two attacks each. This is a cause for concern since it only takes one targeted attack to breach an organisation.

● False positives were not an issue for any of the products.

All of the products allowed all legitimate applications and websites.

● Which products were the most effective?

Products from **Kaspersky**, **Broadcom** and **CrowdStrike** produced extremely good results due to a combination of their ability to block malicious URLs, handle exploits and correctly classify applications and websites. Products from **Trellix**, **Sophos** and **Microsoft** achieved Protection Accuracy scores in the high 90s. All products received AAA awards for achieving excellent Total Accuracy ratings.

Security Evaluation

EPS Protection Enterprise Awards

The following products win SE Labs awards:

Broadcom Symantec Endpoint Security

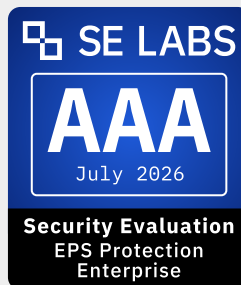
Kaspersky Endpoint Security

CrowdStrike Falcon

Trellix Endpoint Security

Sophos Endpoint Security

Microsoft Defender Antivirus (enterprise)



SE LABS PRESENTS

THE - C2

MARCH 2027

Connecting business with cyber security

The-C2 is an exclusive, invite-only threat intelligence event that connects multinational business executives with the cutting edge of the cyber security industry. The event enables frank and open discussion of the developing digital threat landscape among global security leaders.

The-C2 is hosted by SE Labs, the world's leading security testing lab. Its unique position in the industry provides a route to understanding both the developing threat landscape and the evolving security measures for defending against attackers.

THE - C2 . COM

Threat Responses

Full Attack Chain: Testing Every Layer of Detection and Protection

Attackers start from a certain point and don't stop until they have either achieved their goal or have reached the end of their resources (which could be a deadline or the limit of their abilities). This means that, in a test, the tester needs to begin the attack from a realistic first position, such as sending a phishing email or setting up an infected website, and moving through many of the likely steps leading to actually stealing data or causing some other form of damage to the network.

If the test starts too far into the attack chain, such as executing malware on an endpoint, then many products will be denied opportunities to use the full extent of their protection and detection abilities. If the test concludes before any 'useful'

damage or theft has been achieved, then similarly the product may be denied a chance to demonstrate its abilities in behavioural detection and so on.

Attack Stages

The illustration (below) shows typical stages of an attack. In a test, each of these should be attempted to determine the security solution's effectiveness. This test's results record detection and protection for each of these stages.

We measure how a product responds to the first stages of the attack with a detection and/or protection rating. Sometimes products allow threats to run yet still detect them. Other times they might allow the threat to run briefly before neutralising it.

Ideally, they detect and block the threat before it has a chance to run. Products may delete threats or automatically contain them in a 'quarantine' or other safe holding mechanism for later analysis.

Should the initial attack phase succeed, we then measure post-exploitation stages, which are represented by steps two through to seven below. We broadly categorise these stages as: Access (step 2); Action (step 3); Escalation (step 4); and Post-Escalation (steps 5-6).

In figure 1. you can see a typical attack running from start to end, through various 'hacking' activities. This can be classified as a fully successful breach.

In figure 2. a product or service has interfered with the attack, allowing it to succeed only as far as stage 3, after which it was detected and neutralised. The attacker was unable to progress through stages 4 onwards.

It is possible for an attack to run in a different order with, for example, the attacker attempting to connect to other systems without needing to escalate privileges. However, it is common for password theft (see step 5) to occur before using stolen credentials to move further through the network.

How Hackers Progress

Figure 1. A typical attack starts with an initial contact and progresses through various stages, including reconnaissance, stealing data and causing damage.



Figure 2. This attack was initially successful but only able to progress as far as the reconnaissance phase.



1. Protection and Legitimate Handling Accuracy

1.1 Protection Details

These results break down how each product handled threats into some detail. You can see how many detected a threat and the levels of protection provided.

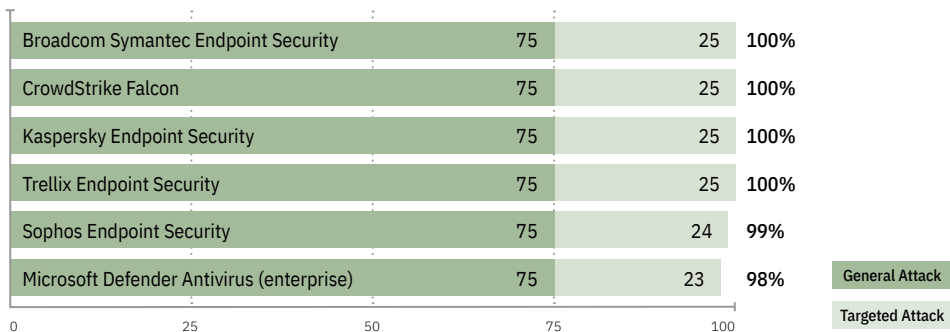
Products sometimes detect more threats than they protect against. This can happen when they recognise an element of the threat but aren't equipped to stop it. Products can also provide protection even if they don't detect certain threats. Some threats abort on detecting specific endpoint protection software.

Product	Detected	Blocked	Neutralised	Compromised	Protected
Broadcom Symantec Endpoint Security	100	100	0	0	100
CrowdStrike Falcon	100	100	0	0	100
Kaspersky Endpoint Security	100	100	0	0	100
Trellix Endpoint Security	100	100	0	0	100
Sophos Endpoint Security	100	99	0	1	99
Microsoft Defender Antivirus (enterprise)	100	98	0	2	98

- This data shows in detail how each product handled the threats used.

1.2 Attack Types

The graph shows how each product protected against the different types of attacks used in the test.

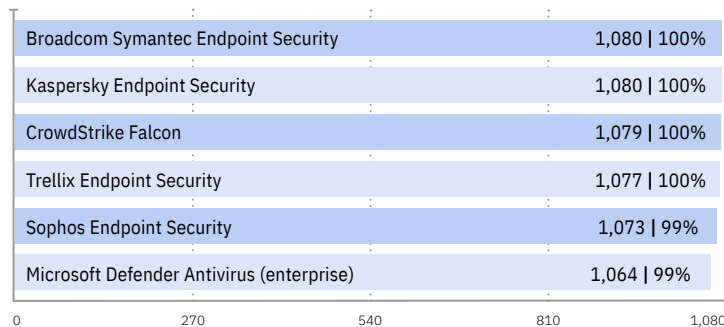


1.3. Total Accuracy Ratings

Judging the effectiveness of an endpoint security product is a subtle art, and many factors are at play when assessing how well it performs. To make things easier we've combined all the different results from this report into one easy-to-understand graph.

The graph takes into account not only each product's ability to detect and protect against threats, but also its handling of non-malicious objects such as web addresses (URLs) and applications.

Not all protections, or detections for that matter, are equal. A product might completely block a URL, which stops the threat before it can even start its intended series of malicious events. Alternatively, the product might allow a web-based exploit to execute but prevent it from downloading any further code to the target.



- Categorising how a product handles legitimate objects is complex, and you can find out how we do it in **Legitimate Accuracy Ratings** on page 10.

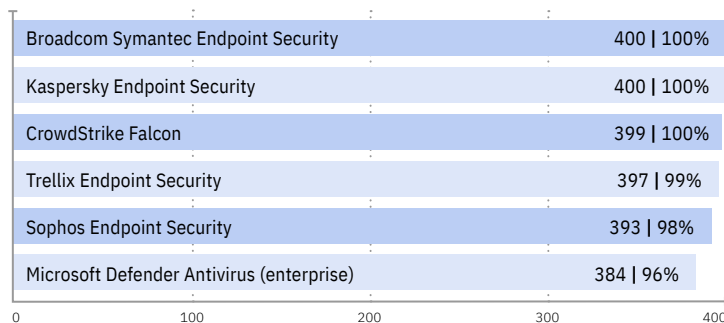
- Total Accuracy Ratings combine protection and false positives

In another case malware might run on the target for a short while before its behaviour is detected and its code is deleted or moved to a safe 'quarantine' area for future analysis. We take these outcomes into account when attributing points that form final ratings.

For example, a product that completely blocks a threat is rated more highly than one that allows a threat to run for a while before eventually evicting it. Products that allow all malware infections, or that block popular legitimate applications, are penalised heavily.

1.4 Protection Accuracy

To understand how we calculate these ratings, see **Appendix A: Protection Ratings** on page 12.



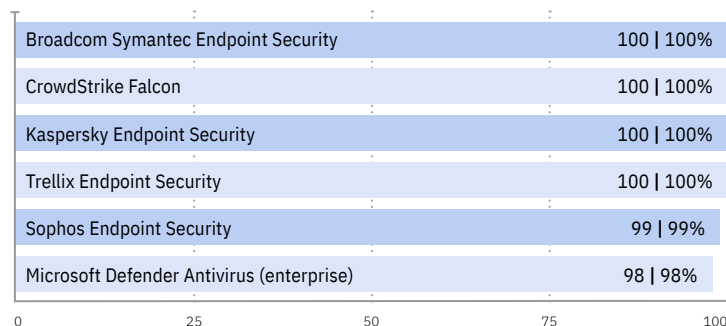
- Protection Accuracy Ratings are weighted to show that how products handle threats can be subtler than just 'win' or 'lose'.

Average 99%

1.5 Protection Scores

This graph shows the overall level of protection, making no distinction between neutralised and blocked incidents.

For each product we add Blocked and Neutralised cases together to make one simple tally.



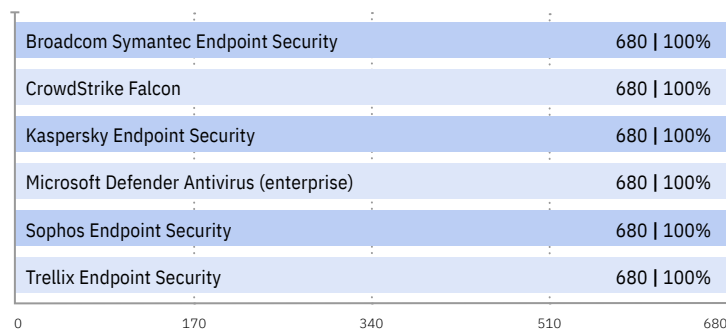
● Protection Scores are a simple count of how many times a product protected the system.

1.6 Legitimate Accuracy Ratings

These ratings indicate how accurately the products classify legitimate applications and URLs, while also taking into account the interactions that each product has with the user. Ideally a product will either not classify a legitimate object or will classify it as safe. In neither case should it bother the user.

We also take into account the prevalence (popularity) of the applications and websites used in this part of the test, applying stricter penalties for when products misclassify very popular software and sites.

To understand how we calculate these ratings, see **Accuracy Ratings** on page 14.



● Legitimate Accuracy Ratings can indicate how well a vendor has tuned its detection engine.

2. Conclusion

Attacks in this test included threats that affect the wider public and more closely targeted individuals and organisations. You could say that we tested the products with ‘public’ malware and full-on hacking attacks.

We introduced the threats in a realistic way such that threats seen in the wild on websites were downloaded from those same websites, while threat caught sending email were delivered to our target systems as emails. While we do ‘create’ threats by using publicly available free hacking tools, we do not write unique malware so there is no technical reason why any vendor being tested should do poorly.

All the products are well-known and should do well in this test. The results were generally strong in the way the products handled both public threats and targeted attacks. All of the six products achieved excellent Protection Accuracy Ratings with half of them stopping all of the attacks.

The products which achieved 100% Protection Accuracy Ratings were **Kaspersky Endpoint Security**, **Broadcom Symantec Endpoint Security** and **CrowdStrike Falcon**. These products detected and protected against all of the public malware

used in this test. These threats are downloaded from the Internet on the day that the products are tested and demonstrate a product’s familiarity with common threats and its ability to keep its databases current.

Indeed, all the products were very capable of handling Web downloads. **Sophos Endpoint** and **Microsoft Defender Antivirus (enterprise)** also repelled all such attacks.

A single point was docked from **Trellix Endpoint Security** for an artifact that remained in the system after a Web attack that was eventually blocked. The product also lost a couple of points for two targeted attacks that ran briefly before they were blocked.

Sophos Endpoint lost three points when it detected, but was unable to block, one targeted attack. **Microsoft Defender Antivirus (enterprise)** also missed two targeted attacks and lost a point for an artifact left by a malicious process that was eventually blocked.

The products did recognise that malicious processes were running and attempted to block them. But targeted attacks are more persistent.

Once the initial malicious file is downloaded and executed, the targeted attack often progresses unimpeded all the way up to the post-escalation stage. This demonstrates that a robust reaction when the targeted attack is already underway is not as effective as an outright block at the onset of an attack. This is how the products from **Kaspersky**, **Broadcom** and **CrowdStrike** successfully defended against all targeted attacks.

All of the products achieved 100% Legitimate Accuracy Ratings and correctly classified legitimate files and websites as safe.

All the products in this test posted excellent results and won AAA awards. The strongest, from **Kaspersky**, **Broadcom** and **CrowdStrike** scored 100% Protection Accuracy Ratings by providing protection against all of the threats. **Trellix Endpoint Security** was only a single percentage point shy of achieving the same, while the enterprise products from **Sophos** and **Microsoft** posted Protection Accuracy Ratings in the high 90s.

Appendices

Appendix A: Protection Ratings

The results below indicate how effectively the products dealt with threats. Points are earned for detecting the threat and for either blocking or neutralising it.

■ Detected (+1) If the product detects the threat with any degree of useful information, we award it one point.

■ Blocked (+2) Threats that are disallowed from even starting their malicious activities are blocked. Blocking products score two points.

■ Complete Remediation (+1) If, in addition to neutralising a threat, the product removes all significant traces of the attack, it gains an additional one point.

■ Neutralised (+1) Products that kill all running malicious processes 'neutralise' the threat and win one point.

■ Persistent Neutralisation (-2) This result occurs when a product continually blocks a persistent threat from achieving its aim, while not removing it from the system.

■ Compromised (-5) If the threat compromises the system, the product loses five points. This loss may

be reduced to four points if it manages to detect the threat (see Detected, above), as this at least alerts the user, who may now take steps to secure the system.

Rating Calculations

We calculate the protection ratings using the following formula:

Protection Rating =
(1x number of Detected) +
(2x number of Blocked) +
(1x number of Neutralised) +
(1x number of Complete Remediation) +
(-5x number of Compromised)

The 'Complete Remediation' number relates to cases of neutralisation in which all significant traces of the attack were removed from the target.

These ratings are based on our opinion of how important these different outcomes are. You may have a different view on how seriously you treat a 'Compromise' or 'Neutralisation without complete remediation'. If you want to create your own rating system, you can use the raw data from **1.1 Protection Details** on page 8 to roll your own set of personalised ratings.

Targeted Attack Scoring

The following scores apply only to targeted attacks and are cumulative, ranging from -1 to -5.

■ Access (-1) If any command that yields information about the target system is successful this score is applied. Examples of successful commands include listing current running processes, exploring the file system and so on. If the first command is attempted and the session is terminated by the product without the command being successful the score of Neutralised (see above) will be applied.

■ Action (-1) If the attacker is able to exfiltrate a document from the target's Desktop of the currently logged in user then an 'action' has been successfully taken.

■ Escalation (-2) The attacker attempts to escalate privileges to NT Authority/System. If successful, an additional two points are deducted.

■ Post-Escalation Action (-1) After escalation the attacker attempts actions that rely on escalated privileges. These include attempting to steal credentials, modifying the file system and recording keystrokes. If any of these actions are successful then a further penalty of one point deduction is applied.

Appendix B: Legitimate Interaction Ratings

It's **crucial** that anti-malware endpoint products not only stop – or at least detect – threats, but that they allow legitimate applications to install and run without misclassifying them as malware. Such an error is known as a 'false positive' (FP).

In reality, genuine FPs are quite rare in testing. In our experience it is unusual for a legitimate application to be classified as 'malware'. More often it will be classified as 'unknown', 'suspicious' or 'unwanted' (or terms that mean much the same thing).

We use a subtle system of rating an endpoint's approach to legitimate objects, which takes into account how it classifies the application and how it presents that information to the user. Sometimes

the endpoint software will pass the buck and demand that the user decide if the application is safe or not. In such cases the product may make a recommendation to allow or block. In other cases, the product will make no recommendation, which is possibly even less helpful.

If a product allows an application to install and run with no user interaction, or with simply a brief notification that the application is likely to be safe, it has achieved an optimum result. Anything else is a Non-Optimal Classification/Action (NOCA). We think that measuring NOCAs is more useful than counting the rarer FPs.

Prevalence Ratings

There is a significant difference between an

Legitimate Software Prevalence Rating Modifiers

Very High Impact	5
High Impact	4
Medium Impact	3
Low Impact	2
Very Low Impact	1

endpoint product blocking a popular application such as the latest version of Microsoft Word and condemning a rare Iranian dating toolbar for Internet Explorer 6. One is very popular all over the world and its detection as malware (or something less serious but still suspicious) is a big deal. Conversely, the outdated toolbar won't have had a comparably large user base even when it was new. Detecting this application as malware may be wrong, but it is less impactful in the overall scheme of things.

With this in mind, we collected applications of varying popularity and sorted them into five separate categories, as follows:

1. Very High Impact
2. High Impact
3. Medium Impact
4. Low Impact
5. Very Low Impact

	None (allowed)	Click to Allow (default allow)	Click to Allow/ Block (no recommendation)	Click to Block (default block)	None (blocked)	
Safe	2	1.5	1			A
Unknown	2	1	0.5	0	-0.5	B
Not Classified	2	0.5	0	-0.5	-1	C
Suspicious	0.5	0	-0.5	-1	-1.5	D
Unwanted	0	-0.5	1	-1.5	-2	E
Malicious				2	-2	F
	1	2	3	4	5	

Incorrectly handling any legitimate application will invoke penalties, but classifying Microsoft Word as malware and blocking it without any way for the user to override this will bring far greater penalties than doing the same for an ancient niche toolbar. In order to calculate these relative penalties, we assigned each impact category with a rating modifier, as shown in the table above.

Applications were downloaded and installed during the test, but third-party download sites were avoided and original developers' URLs were used where possible. Download sites will sometimes bundle additional components into applications' install files, which may correctly cause anti-malware products to flag adware. We remove

adware from the test set because it is often unclear how desirable this type of code is.

The prevalence for each application and URL is estimated using metrics such as third-party download sites and the data from Tranco.com's global traffic ranking system.

Accuracy Ratings

We calculate legitimate software accuracy ratings by multiplying together the interaction and prevalence ratings for each download and installation:

Accuracy rating = Interaction rating x Prevalence rating

If a product allowed one legitimate, Medium impact application to install with zero interaction with the user, then its Accuracy rating would be calculated like this:

Accuracy rating = 2 x 3 = 6

This same calculation is made for each legitimate application/site in the test and the results are summed and used to populate the graph and table shown under **Legitimate Accuracy Ratings** on page 10.

Distribution of Impact Categories

Endpoint products that were most accurate in handling legitimate objects achieved the highest ratings. If all objects were of the highest prevalence, the maximum possible rating would be 1,000 (100 incidents x (2 interaction rating x 5 prevalence rating)).

In this test there was a range of applications with different levels of prevalence. The table below shows the frequency:

Legitimate Software Category Frequency

Prevalence Rating	Frequency
High Impact	40
Medium Impact	40
Low Impact	20

Legitimate Interaction Ratings

Product	None (allowed)	Click to allow/block (No Recommendation)	None (blocked)
Broadcom Symantec Endpoint Security	100	0	0
CrowdStrike Falcon	100	0	0
ESET Endpoint Security	100	0	0
Kaspersky Endpoint Security	100	0	0
Microsoft Defender Antivirus (enterprise)	100	0	0
Sophos Endpoint Security	100	0	0

● Products that do not bother users and classify most applications correctly earn more points than those that ask questions and condemn legitimate applications.

Appendix C: Terms Used

Compromised The attack succeeded, resulting in malware running unhindered on the target. In the case of a targeted attack, the attacker was able to take remote control of the system and carry out a variety of tasks without hindrance.

Blocked The attack was prevented from making any changes to the target.

False Positive When a security product misclassifies a legitimate application or website as being malicious, it generates a 'false positive'.

Neutralised The exploit or malware payload ran on the target but was subsequently removed.

Complete Remediation If a security product removes all significant traces of an attack, it has achieved complete remediation.

Target The test system that is protected by a security product.

Threat A program or sequence of interactions with the target that is designed to take some level of unauthorised control of that target.

Update Security vendors provide information to their products in an effort to keep abreast of the latest threats. These updates may be downloaded in bulk as one or more files, or requested individually and live over the internet.

Appendix D: FAQs

Q What is a partner organisation? Can I become one to gain access to the threat data used in your tests?

A Partner organisations benefit from our consultancy services after a test has been run. Partners may gain access to low-level data that can be useful in product improvement initiatives and have permission to use award logos, where appropriate, for marketing purposes. We do not share data on one partner with other partners. We do not partner with organisations that do not engage in our testing.

Q I am a security vendor and you tested my product without permission. May I access the threat data to verify that your results are accurate?

A We are willing to share a certain level of test data with non-partner participants for free. The intention is to provide sufficient data to demonstrate that the results are accurate. For more in-depth data suitable for product improvement purposes we recommend becoming a partner.

A full methodology for this test is available from our website.

- The test was conducted between 6th April and 12th June 2026.
- All products were configured according to each vendor's recommendations, when such recommendations were provided.
- Targeted attacks were selected and verified by SE Labs.
- Malicious emails, URLs, attachments and legitimate messages were independently located and verified by SE Labs.
- Malicious and legitimate data was provided to partner organisations once the test was complete.

Appendix E: Product Versions

The table below shows the service's name as it was being marketed at the time of the test.

Vendor	Product	Build Version (start)	Build Version (end)
Broadcom	Symantec Endpoint Security	SEP Version: 16.0.0	SEP Version: 16.0.0
CrowdStrike	Falcon	7.35.20709.0	7.39.21105.0
Kaspersky	Endpoint Security	12.11.0.637	12.11.0.637
Microsoft	Defender Antivirus (enterprise)	Antimalware Client Version: 4.18.26020.6 Engine Version: 1.1.26020.3 Antivirus Version: 1.447.102.0 Anti-spyware Version: 1.447.102.0	Antimalware Client Version: 4.18.26050.15 Engine Version: 1.1.26050.11 Antivirus Version: 1.453.257.0 Anti-spyware Version: 1.453.257.0
Sophos	Endpoint Security	Core Agent: 2025.2.3.8.0 Sophos Intercept: 2024.1.2.1.0 Device Encryption: 2025.1.3.2.0	Core Agent: 2026.2.0.42.0 Sophos Intercept: 2024.1.2.1.0 Device Encryption: 2025.1.3.2.0
Trellix	Endpoint Security	Endpoint Security Platform Version: 10.7.20.16072 Adaptive Threat Protection Version: 10.7.20.13219 Threat Prevention Version: 10.7.20.14066 Firewall Version: 10.7.20.14030 Web Control Version: 10.7.20.13154	Endpoint Security Platform Version: 10.7.20.16072 Adaptive Threat Protection Version: 10.7.20.13219 Threat Prevention Version: 10.7.20.14066 Firewall Version: 10.7.20.14030 Web Control Version: 10.7.20.13154

SE Labs Report Disclaimer

1. The information contained in this report is subject to change and revision by SE Labs without notice.
2. SE Labs is under no obligation to update this report at any time.
3. SE Labs believes that the information contained within this report is accurate and reliable at the time of its publication, which can be found at the bottom of the contents page, but SE Labs does not guarantee this in any way.
4. All use of and any reliance on this report, or any information contained within this report, is solely at your own risk. SE Labs shall not be liable or responsible for any loss of profit (whether incurred directly or indirectly), any loss of goodwill or business reputation, any loss of data suffered, pure economic loss, cost of procurement of substitute goods or services, or other intangible loss, or any indirect, incidental, special or consequential loss, costs, damages, charges or expenses or exemplary damages arising from this report in any way whatsoever.
5. The contents of this report does not constitute a recommendation, guarantee, endorsement or otherwise of any of the products listed, mentioned or tested.
6. The testing and subsequent results do not guarantee that there are no errors in the products, or that you will achieve the same or similar results. SE Labs does not guarantee in any way that the products will meet your expectations, requirements, specifications or needs.
7. Any trade marks, trade names, logos or images used in this report are the trade marks, trade names, logos or images of their respective owners.
8. The contents of this report are provided on an "AS IS" basis and accordingly SE Labs does not make any express or implied warranty or representation concerning its accuracy or completeness.