



脅威の調査

脅威の調査は、脅威および悪意のあるアクティビティを検出するために検査エンジンに適用される一連のルールから生成されます。このページでは、これらのルールを表示できます。

Multicloud Defense は毎日 1 回ネットワーク侵入の新しいルールまたは変更されたルールがないか検索し、ルールを内部ライブラリに含めたり、既知の悪意のあるソースを内部ライブラリから削除したりします。このアクションは自動化されています。この機能には、送信元としての IP アドレスの新しいリストをダウンロードして検証し、それらを新しいルールセットに実装する動作が含まれています。その後、これらのルールセットが展開されます。

ルールには、ポリシー、クラス、アプリケーション、ルールセットライブラリの日付、その他のパラメータなど、さまざまな編成方法があります。阻止されたルール（脅威や悪意のあるアクティビティが検出されたなど）について詳しく知りたい場合は、[脅威の調査（Threat Research）] ページを使用してルールの詳細を確認してください。各ページの次の部分を利用できます。

検索バー

ウィンドウの上部にある検索バーを使用すると、脅威の調査の各ページで、既知の IP アドレス、アクション、ルール名、ゲートウェイ名、攻撃タイプ、またはプロファイル名などの単一の識別ファセットを検索できます。スクロールして特定のフィールド値を見つけた場合は、[検索に追加（Add to Search）] を使用して、検索エクスペリエンスを容易にすることができます。

検索はページごとに独立しているため、さまざまなタイプの脅威の調査を横断的に検索することはできないことに注意してください。詳細については、次のセクションを参照してください。

[詳細の表示（View Details）]

脅威の調査の各ファセットには、単一のインシデントまたは攻撃の [詳細の表示（View Details）] 機能があります。これらの詳細で提供される値は、脅威の調査のタイプによって異なりますが、ポリシー、セキュリティプロファイル、ルール、またはルールセットを微調整する場合に役立ちます。

検索に追加

ここで使用可能な任意のタイプの調査について、行内の任意のいずれか1つの値をクリックすると、[検索に追加 (Add to Search)] オプションが自動的に表示されます。これにより、選択した値がウィンドウの上部にある検索バーに自動的に適用され、検索バーのコンテンツに合わせて表示ウィンドウがフィルタリングされます。これを複数回実行すると、選択した値が合成されて複雑な検索リクエストが生成されます。

- [ネットワーク侵入 \(2 ページ\)](#)
- [Web の保護 \(3 ページ\)](#)
- [悪意のある送信元 \(3 ページ\)](#)

ネットワーク侵入

ネットワーク侵入とは、ネットワーク上での不正なアクティビティを指します。この表には、IDS/IPS エンジンの組み込みルールや、これらのルールの関連情報は含まれていません。これらのルールは検出のみを目的としています。残りの IDS/IPS ルールは、侵入または攻撃のさまざまなレベルに基づいてアクションを実行し、保護するように設定されます。

[ネットワーク侵入 (Network Intrusion)] ページには、次の情報が表示されます。

- [ゲートウェイ名 (Gateway Names)] : 悪意のある送信元を処理した、影響を受けるゲートウェイの名前。
- [プロファイル名 (Profile Names)] : 悪意のある送信元によってトリガーされたセキュリティプロファイルの名前。
- [IPS ポリシー (IPS Policy)] : イベントまたは攻撃によってトリガーされた Multicloud Defense 内のポリシー。
- [IPS クラス (IPS Class)] : 攻撃シグニチャートラフィックのデータベースによって決定される攻撃のタイプ。
- [IPS カテゴリ (IPS Category)] : イベントまたは攻撃によってトリガーされた IPS 署名カテゴリ。
- [ルールID (Rule ID)] : イベントまたは攻撃によってトリガーされた、Multicloud Defense 内で内部的に記載されているルール ID。
- [影響を受けるサービス (Services Impacted)] : イベントまたは攻撃の影響を受ける Web サービスのタイプ。
- [影響 (Impact)] : イベントまたは攻撃による既知または想定される影響の重大度レベル。
- [メッセージ (Message)] : 攻撃として識別されたイベントの内容。
- [ルールコンテンツ (Rule Content)] : イベントまたは攻撃によってトリガーされたルールのコンテンツ。

- [CVSSスコア (CVSS Score)] : 共通脆弱性評価システム (CVSS) は、情報セキュリティの脆弱性の重大度に数値スコアを割り当てるフレームワークです。CVSS スコアの範囲は 0 ～ 10 で、10 が最も重大です。
- [CVE] : Common Vulnerabilities and Exposures (CVE) は、脆弱性を分類する用語一覧です。攻撃またはイベントのタイプに関連付けられている CVE がある場合、内部ライブラリはここでその値を自動的に生成します。
- [参考資料 (References)] : 公開されている場合、このリンクをクリックすると、CVE のアナウンスの原文および分類に移動します。

Web の保護

Web アプリケーション ファイアウォール (WAF) の調査は「Web の保護」として表示されます。この機能を使用すると、Web の脅威からデバイスを保護し、不要なコンテンツを規制できます。[Web の保護 (Web Protection)] ページには、次の情報が表示されます。

- [ゲートウェイ名 (Gateway Names)] : 悪意のある送信元を処理した、影響を受けるゲートウェイの名前。
- [プロファイル名 (Profile Names)] : 悪意のある送信元によってトリガーされたセキュリティプロファイルの名前。
- [CRSカテゴリ (CRS Category)] : 一般的な攻撃検出ルールセットごとに識別されるコアルールセット (CRS) のカテゴリ。
- [インスペクションタイプ (Inspection Type)] : 攻撃またはイベントをカプセル化したトラフィックで Multicloud Defense によって実行されたインスペクションのタイプ。
- [攻撃タイプ (Attack Type)] : ネットワークを通過した不正な攻撃のタイプ。
- [プラットフォーム (Platform)] : 攻撃またはイベントから特定されたプラットフォームのタイプ。
- [言語 (Language)] : イベントから検出された Web 開発言語。

悪意のある送信元

悪意のある送信元とは、ネットワークに損害を与えるあらゆるタイプのコードまたはパケットを指します。[悪意のある送信元 (Malicious Sources)] ページには、次の情報が表示されます。

- [ゲートウェイ名 (Gateway Names)] : 悪意のある送信元を処理した、影響を受けるゲートウェイの名前。
- [プロファイル名 (Profile Names)] : 悪意のある送信元によってトリガーされたセキュリティプロファイルの名前。

- [悪意のある送信元のアクション (Malicious Sources Action)] : 悪意のある送信元が特定されたときに実行されるアクション。
- [影響 (Impact)] : ライブラリ内でのランク付けによって判断される、悪意のあるコンテンツの影響。
- [悪意のある送信元のIP (Malicious Source IP)] : 悪意のある送信元の IP アドレス。

翻訳について

このドキュメントは、米国シスコ発行ドキュメントの参考和訳です。リンク情報につきましては、日本語版掲載時点で、英語版にアップデートがあり、リンク先のページが移動/変更されている場合がありますことをご了承ください。あくまでも参考和訳となりますので、正式な内容については米国サイトのドキュメントを参照ください。