

脅威ハンティング実践のすゝめ

独立行政法人情報処理推進機構 産業サイバーセキュリティセンター

中核人材育成プログラム8期生 脅威インテリジェンスを活用したセキュリティ強化のためのアプローチ

目次

第1章 はじめに	3
1.1 本レポート作成の背景	3
1.2 本レポートの目的	3
1.3 本レポートの構成と対象読者	4
1.3.1 本レポートの全体構成と各章の概要	4
1.3.2 想定読者	4
1.3.3 本レポートから得られる知見	4
1.3.4 免責事項	4
第2章 アクティブサイバーディフェンスとは	5
2.1 アクティブサイバーディフェンスの概念	5
2.2 アクティブサイバーディフェンスの位置づけと境界	6
2.3 国や組織によるアクティブサイバーディフェンスに関する取り組み	9
2.4 日本政府における「能動的サイバー防御」の考え方	11
第3章 脅威インテリジェンスの概要	13
3.1 脅威インテリジェンスとは	13
3.1.1 脅威インテリジェンスの定義	13
3.1.2 脅威インテリジェンスの必要性	13
3.1.3 脅威インテリジェンスのプロセス	15
3.1.4 脅威インテリジェンスを利用する為の前提事項	16
3.2 脅威インテリジェンスのライフサイクル	16
3.2.1 方針策定	16
3.2.2 収集・加工	17
3.2.3 分析	18
3.2.4 配布	19
3.2.5 評価・改善	20
3.3 脅威ハンティングとの関連性	21
第4章 脅威ハンティングの概要	22
4.1 脅威ハンティングとは	22
4.1.1 脅威ハンティングの定義	22
4.1.2 脅威ハンティングの目的と重要性	22
4.1.3 脅威ハンティングを実施する為の前提事項	24
4.2 脅威ハンティングの主なアプローチ	28
4.2.1 攻撃主導型 (Attack based Hunting)	28
4.2.2 データ主導型 (Data based Hunting)	29
4.2.3 エンティティ主導型 (Entity based Hunting)	29
4.2.4 インテリジェンス主導型 (Intel based Hunting)	30
4.2.5 ハイブリッド型 (Hybrid Hunting)	31

4.2.6 構造化・非構造化の視点によるアプローチ整理	31
4.3 脅威ハンティングのライフサイクル	32
4.3.1 計画	33
4.3.2 実行	35
4.3.3 完了	40
第5章 脅威ハンティング実践検証	42
5.1 検証設計	42
5.1.1 検証の目的	42
5.1.2 検証における仮説と前提	42
5.1.3 検証環境	42
5.2 攻撃主導型の脅威ハンティング検証	43
5.2.1 脅威シナリオ	43
5.2.2 実践プロセス	43
5.2.3 検証結果	46
5.2.4 検証結果に対する考察	46
5.3 ハイブリッド型の脅威ハンティング検証	47
5.3.1 脅威シナリオ	47
5.3.2 実践プロセス	47
5.3.3 検証結果	48
5.3.4 検証結果に対する考察	48
6章 まとめ	50
6.1 組織における実践に向けた考察	50
6.1.1 検証から見えた有効性と実践上の課題	50
6.1.2 組織への導入・定着に向けた要点	51
6.2 おわりに	56
謝辞	58
用語集	59

第1章 はじめに

1.1 本レポート作成の背景

近年、企業活動を支える情報システムは、クラウド、IoT、ビッグデータ解析やAIなどの先進技術の導入によって急速に高度化・複雑化している。このような技術革新は業務効率や競争力の向上に寄与する一方で、情報システムが担う役割の重要性を飛躍的に高めており、その依存度も年々増加している。企業にとって、IT インフラは単なる業務支援のツールではなく、経営の根幹を成す存在となっていると言える。

しかしながら、こうしたデジタル環境の拡大は同時に、サイバー攻撃者にとって新たな攻撃対象や侵入経路を提供する要因ともなっている。実際に、標的型攻撃、ランサムウェア、サプライチェーン攻撃、ゼロデイ攻撃を始めとした攻撃手法は高度化・巧妙化しており、企業システムは依然としてサイバー空間における魅力的な攻撃対象であることに変わりはない。特に、正規の通信を装ってネットワークをすり抜けたり、既存の防御システムを回避して長期間潜伏するようなステルス型の手法は、従来の検知技術では見逃されるリスクが高い。

従来のセキュリティ対策は、主に境界防御やシグネチャベースの検知といった「防御」および「検知」に重点を置いた受動的な手法で構成されており、未知の攻撃や既存パターンに該当しない異常を早期に検知することは困難である。結果として、攻撃の兆候を見逃し、インシデントが発生してから初めて対応に追われる「後手の対応」が多くの現場で課題となっている。

このような現状を踏まえ、注目されているのが「アクティブサイバーディフェンス」と呼ばれるセキュリティ対策の概念である。なかでも「脅威ハンティング」は、防御側が能動的にシステム内部の挙動を分析し、セキュリティ製品のアラートや既知のパターンに依存することなく、潜在的な脅威を探索・特定する取り組みであり、未知の攻撃に対応するための有効な手段とされている。

1.2 本レポートの目的

脅威ハンティングは、既存の防御システムによる検知に頼らず、組織内に潜む未検出の脅威を能動的に探索するアプローチであり、従来の受動的な防御モデルでは見逃されがちな高度な脅威を早期に発見し、インシデントの被害を未然に防ぐことが可能となる。ただし、脅威ハンティングの導入・運用には、限られたリソースの中でどのように優先順位を付け、効果的に運用していくかという課題が浮き彫りになる。実際に多くの組織では、データの収集や分析能力、専門人材の不足といった課題に直面しており、脅威ハンティングの導入に際しては組織全体の体制づくりも重要なポイントとなる。本レポートは、脅威ハンティングの手法やプロセスを整理するとともに、実践的な検証を通じて得られた導入上の課題や現実的なアプローチに関する考察を提供することを目的としている。

1.3 本レポートの構成と対象読者

1.3.1 本レポートの全体構成と各章の概要

本レポートは、全6章で構成されており、アクティブサイバーディフェンスの概念および脅威ハンティングの基礎的理論を解説するとともに、模擬環境における検証を踏まえた脅威ハンティング実践のための考察について説明する。

各章の概要は以下の通りである。

第2章：アクティブサイバーディフェンスの概念や注目される背景について解説する。

第3章：脅威インテリジェンスの収集・分析・活用方法、導入・運用における考慮事項などを解説する。

第4章：脅威ハンティングの概念、必要性、実施手順、考慮事項などを解説する。

第5章：脅威ハンティングの実践的な有効性を検証するために、本プロジェクトで実施した検証内容と得られた知見・考察を提示する。

第6章：本プロジェクトを通じて見えた脅威ハンティングの有効性や課題における考察と、組織への導入・定着に向けた要点を解説する。

1.3.2 想定読者

本レポートは、サイバーセキュリティに関する基本的な知識を有し、組織内で脅威の監視に関わる技術者を主な読者として想定している。また、組織のセキュリティ成熟度に課題を感じていたり、既存の監視運用の見直しを検討している方々に向けて、参考となる情報を提供することを意図している。

1.3.3 本レポートから得られる知見

本レポートでは、脅威ハンティングの概要や基本的なアプローチ、プロセスについて整理し、初めて取り組む方にも理解しやすい形で基礎を解説している。また、実践的な検証により得られた知見をもとに、導入時に直面しやすい課題などについても考察している。これにより、限られたリソースの中でも現実的に実施可能な取り組み方や、脅威ハンティングを効果的に運用するための組織的な視点、日常の業務に取り入れるための工夫など、実践に向けた第一歩となる知見を提供する。

1.3.4 免責事項

- 本レポートは単に情報として提供され、内容は予告なしに変更される場合がある。
- 本レポートに誤りがないことの保証や、商品性または特定目的への適合性の黙示的な保証や条件を含め明示的 または黙示的な保証や条件は一切ないものとする。
- 本レポートに記載の内容は、独立行政法人 情報処理推進機構および産業サイバーセキュリティセンターの意見を代表するものではなく、著者の見解に基づいている。
- 本レポートの利用によるトラブルに対し、本レポート著者ならびに監修者は一切の責任を負わないものとする。
- 本レポートの有効期限は、発行日から2年間とする。

第2章 アクティブサイバーディフェンスとは

本章では、アクティブサイバーディフェンスの基本的な概念を説明するとともに、それが従来のパッシブディフェンスとどのように異なるのか解説する。続いて、主要国や国際機関などがこの概念をどのように捉えているのかを比較しながら、国際的な定義の差異について説明する。最後に、日本政府の「能動的サイバー防御」に対する考え方や今後進められる取り組み等について述べる。

2.1 アクティブサイバーディフェンスの概念

アクティブサイバーディフェンス（Active Cyber Defense）は、近年のサイバーセキュリティ分野において注目されている概念であるが、現時点では国際的に統一された定義は存在していない。各国政府、国際機関、研究機関などがそれぞれの立場や目的に応じて独自に定義を行っており、その範囲や適用方針には大きな幅がある。

ただし、多くの定義に共通して見られる考え方として、「従来の受動的な対策と異なり、脅威に対して積極的かつ予防的に対応する姿勢」という点が挙げられる。このようなアプローチは、従来の受動的な防御策では対応が難しい高度なサイバー攻撃に対抗するために必要とされており、攻撃の早期発見と迅速な対応を可能にすることで、国家や組織のセキュリティを強化することが期待されている。

例えば、アメリカのサイバーセキュリティ及び国土安全保障に関連する政策、研修、教育を行う Center for Cyber & Homeland Security（CCHS）の報告書”INTO THE GRAY ZONE: The Private Sector and Active Defense Against Cyber Threats”¹（2016年）によるとアクティブ（サイバー）ディフェンスを「防御的な行動であり、攻撃的な行動ではない」と明確に定義している。また、民間組織がアクティブディフェンスを実施するための具体的な手段として、以下の方法が提案されている。

① 攻撃の検出と分析

- センサーやソフトウェアを用いて、攻撃の兆候を早期に検出すること
- 攻撃の性質や出所を分析し、対応策を講じること

② 攻撃の阻止と緩和

- 攻撃をリアルタイムで阻止するための技術的手段を講じること
- 攻撃の影響を最小限に抑えるための対策を実施すること

③ 情報の収集と共有

- 攻撃者の行動に関する情報を収集し、分析すること
- 得られた情報を適切な関係者と共有し、対応策を協議すること

¹ 「INTO THE GRAY ZONE: The Private Sector and Active Defense Against Cyber Threats」
<https://wayback.archive-it.org/5184/20190103002934/https://cchs.gwu.edu/sites/g/files/zaxdzs2371/f/downloads/CCHS-ActiveDefenseReportFINAL.pdf>

これらの活動は、民間企業が自らのネットワークやシステムを保護するために実施するものであり、政府の法的枠組みや政策の中で行われるべきであるとされている。

次に、Cooperative Cyber Defence Centre of Excellence (CCDCOE) が公表した報告書”Comparative Study on the Cyber Defence of NATO Member States”² (2022 年) によると、「アクティブサイバーディフェンス」という用語は、使用される文脈、すなわち軍事（国家）と民間（非軍事）によってその意味や運用が異なるとされている。一部の専門家は、この違いは実質的な内容よりも、「防衛」という概念が国家の根本的機能として文化的・法的にどう理解されているかによるものだと指摘している。以下は、これらを踏まえてアクティブサイバーディフェンスに関する考え方を国家の安全保障上の観点と民間組織のサイバーセキュリティ上の観点に分けて整理したものである。

① 国家の安全保障上の観点

国家によるアクティブサイバーディフェンスは、憲法、国際法、国内法の枠組みに則り、安全保障政策や外交政策の一環として慎重に検討される。そのため、軍や防衛機関が関与するケースでは、民間組織と比較して、より広い権限のもとで積極的な措置を講じる可能性があると考えられている。

② 民間組織のサイバーセキュリティ上の観点

民間組織におけるアクティブサイバーディフェンスは、攻撃の早期発見や影響の最小化を目的とした戦術的対応とされる。ここでは、国家主権の行使や法的権限の問題が絡むため、軍のように積極的な対抗措置を取ることは制限されるが、内部ネットワーク内での脅威ハンティングやサイバー欺瞞 (deception) など、合法的な範囲での予防的措置が導入されている。

これらの観点から、アクティブサイバーディフェンスとは、定義する主体によって性質自体が変わるのではなく、それを実行する主体の法的・制度的立場に応じて、許容される手段や対応範囲が異なるという点に本質的な違いがあると考えられる。

いずれにしても、アクティブサイバーディフェンスは、単なる防御にとどまらず、脅威の発見・分析・抑止といった一連の積極的対応を通じて、組織や国家の即応力およびレジリエンスを高めることを目的としており、今後のサイバーセキュリティ戦略において重要な要素として位置づけられている。

2.2 アクティブサイバーディフェンスの位置づけと境界

前節でも紹介したアメリカのサイバーセキュリティ及び国土安全保障に関連する政策、研修、教育を行う Center for Cyber & Homeland Security (CCHS) の報告書”INTO THE GRAY ZONE: The Private Sector and Active Defense Against Cyber Threats”³ (2016 年) では、アクティブサイバーディフェンスを

² 「Comparative Study on the Cyber Defence of NATO Member States」 [Comparative-study-on-the-cyber-defence-of-NATO-Member-States.pdf](https://ccdcOE.org/publications/comparative-study-on-the-cyber-defence-of-nato-member-states.pdf)

³ 「INTO THE GRAY ZONE: The Private Sector and Active Defense Against Cyber Threats」
<https://wayback.archive-it.org/5184/20190103002934/https://cchs.gwu.edu/sites/g/files/zaxdzs2371/f/downloads/CCHS-ActiveDefenseReportFINAL.pdf>

正確に理解するためには、その範囲（上限と下限）を明確にする必要があると説明されており、図1に示すとおりである。この図では、相対的な影響とリスクに基づきアクティブ（サイバー）ディフェンスの手段が左から右にランク付けして並べられている。

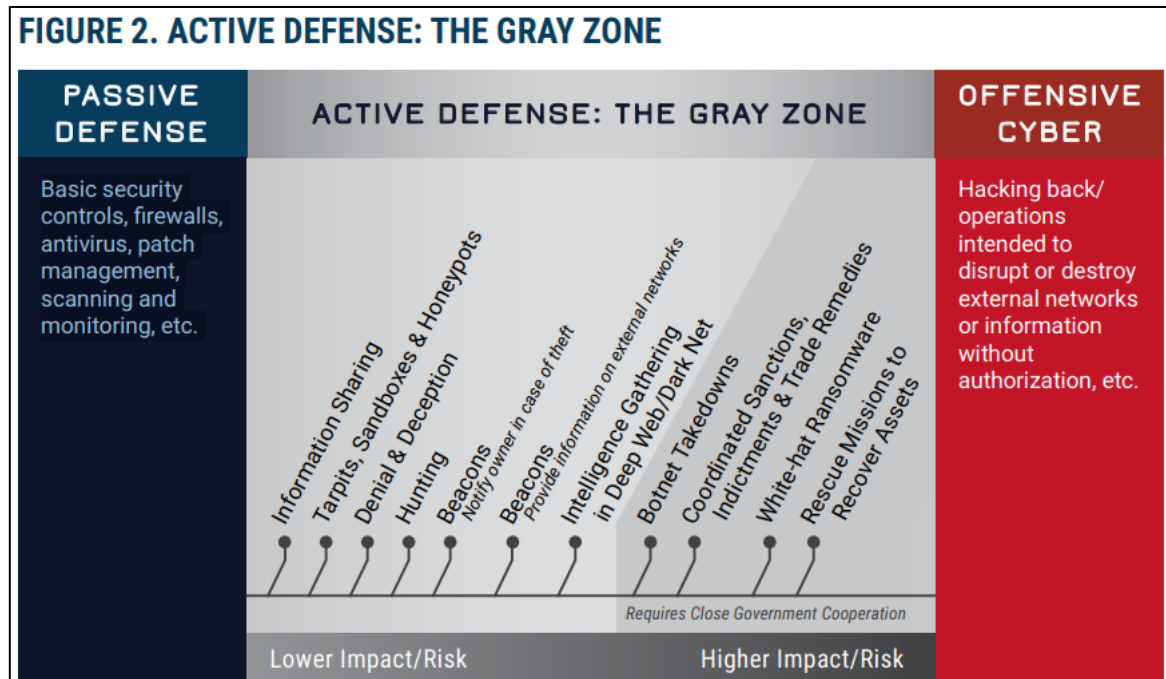


図 1 アクティブ（サイバー）ディフェンスの範囲

（出典：「INTO THE GRAY ZONE: The Private Sector and Active Defense Against Cyber Threats」）

図1で示されているように、サイバーセキュリティにおける活動は、大きく3つの領域に分類される。以下では、それぞれの領域の特徴と代表的な手法について説明する。

① パッシブディフェンス（自分自身のネットワーク内で完結する活動）

自社ネットワーク内で完結する「受動的防御（パッシブディフェンス）」は、外部からの脅威に対する基本的な対策で、例としてはファイアウォール、ウイルス対策ソフト、OSやアプリのパッチ管理などが挙げられる。これらは通常、攻撃の兆候を能動的に探索したり、リアルタイムで戦術的に対処したりするのではなく、あらかじめ設定されたルールに基づいて防御を行う点で、「受動的」と位置づけられる。

さらに、ホワイトリスト（許可されたアプリやIPアドレスの一覧）やブラックリスト（遮断対象の一覧）の管理、管理者権限の制限といった手順的な対応も受動的防御に含まれる。これらは組織のサイバーセキュリティの基本であり、必要不可欠である。しかし、今日の高度化するサイバー脅威には、これだけでは十分に対処できないのが現実である。

② オフェンシブ・サイバー（自分のネットワーク外に影響を及ぼす活動）

「攻撃的サイバー（オフェンシブ・サイバー）」は、パッシブディフェンスとは対極の概念であり、自社ネットワークの外部に対して影響を及ぼす活動を含む。主な目的は、攻撃者に対して損害を与えることである。例えば、攻撃者が盗んだデータを回収または削除する行為や、攻撃者の手法や使用しているツール、意図などの情報を収集する「ハックバック（hack back）」がこれに該当する。

さらに、攻撃者に損害を与えるためのマルウェアの使用、報復的な DDoS 攻撃、知的財産を搾取するために相手システムを積極的に侵害する行為なども、攻撃的サイバーの範疇に入る。

これらの行為は極めて危険であり、民間企業が独自に実施すべきではない。国家安全保障の観点からも、こうした攻撃的手法は、原則として政府主導または明確な法的委任がある場合に限られるべきである。

③ アクティブ（サイバー）ディフェンス（中間に位置する活動）

アクティブサイバーディフェンスは、「自己のシステムの安全確保や運用の自由の維持を目的とする限り、防御的な性質を持つ」とされる。この中間領域、すなわちグレーゾーンには、攻撃者の行動観察を目的としたハニーポットや欺瞞技術、ネットワーク内部に潜伏した脅威の排除（脅威ハンティング）、データ流出先を特定するビーコンの利用、ダークウェブでのインテリジェンス収集などが含まれる。（表 1 の No.1～7）

さらに、ボットネットのトラフィックのシンクホール化、あるいはそのインフラ自体を無力化するような活動も存在するが、これらは「最も攻撃的であり、民間企業は政府と密接に協力する場合にのみ実施すべき」と明示されている。（表 1 の No.8～11）

重要なのは、これらの手法が防御を目的としており、攻撃者への報復や損害を与える行為ではないという点である。つまり、アクティブサイバーディフェンスは「積極的な防御」ではあるものの、「攻撃」ではないと整理されている。

表 1 アクティブサイバーディフェンス（グレーゾーン）

	項目	内容
No.1	情報共有	防御側間で実用的なサイバー脅威指標、緩和ツール、レジリエンス戦略を共有することで、広範な状況認識と防御能力を向上させる。
No.2	ターピット、サンドボックス、ハニーポット	ハッカーのネットワーク境界での活動を阻止する技術ツール、隔離されたオペレーティングシステム内の信頼できないコードの正当性をテストする技術ツール、罠としてセグメント化されたサーバーにハッカーを誘い込み、監視することでハッカーの行動に関する情報を収集する技術ツール。
No.3	否認と欺瞞	偽の情報を混ぜることで、攻撃者が正当な情報に確実にアクセスすることを阻止し、攻撃者に疑念を抱かせ、混乱を招き入れる。
No.4	ハンティング	受動的な防御を回避した後、防御側のネットワーク内に存在する攻撃者を迅速に検出し、的確に排除するための手順と技術対策。
No.5	ビーコン（通知）	ファイルに隠されたソフトウェアまたはリンク。権限のないユーザーがファイルをホームネットワークから削除しようとした場合に、防御者に警告を送信する。
No.6	ビーコン（情報）	ファイルに隠されたソフトウェアまたはリンク。権限なくシステムから削除された場合、通過する外部コンピュータシステムの構造と場所に関する詳細情報を含む接続を防御者に送信できる。

No.7	ディープウェブ・ ダークネットにお ける情報収集	悪意のあるサイバーアクターが通常関心を示すインターネット領域において、秘密裏に監視、なりすまし、資産の偽装などの人的情報技術を用いて、ハッカーの動機、活動、能力に関する情報を取得する活動。
No.8	ボットネットのテ イクダウン	侵害されたコンピュータネットワークの C2 インフラストラクチャから、マルウェアに感染した多数のコンピュータを特定し、切断する技術的アクション。
No.9	協調的な制裁、起 訴、貿易救済措置	既知の悪意あるサイバー行為者に対し、資産の凍結、法的訴追、そして行為者やその国家スポンサーを標的とした懲罰的貿易政策の実施などを通じて、民間部門と政府が協調してコストを課す行動。
No.10	ホホワイトハットラ ンサムウェア	第三者のコンピュータシステム上のファイルを、マルウェアを用いて合法的に暗号化する。このファイルには、悪意ある行為者のシステムへ転送中の盗まれた情報が含まれる。その後、官民連携のパートナーは、被害を受けた第三者に対し、侵害を受け、盗まれた資産を所有していることを通知する。ファイルへのアクセスを回復するには、盗まれた資産を返却する必要がある。
No.11	資産回収のための レスキューミッシ ョン	ハッキングツールを使用して、情報を盗んだ敵対者のコンピュータネットワークに侵入し、その情報の侵害の程度を特定し、最終的に回復しようとする。

このように、アクティブサイバーディフェンスは、従来の受動的防御では対応が難しい高度なサイバー脅威に対処するための手段である。ただし、攻撃的手法とならないよう慎重な設計が必要であり、民間企業が実施する際には、法的・倫理的な枠組みを守り、政府との連携や規制を踏まえた対応が求められる。アクティブとパッシブの両防御は対立するものではなく、相互に補完する関係にある。現代の複雑なサイバー攻撃に対抗するには、受動的防御に加えて、アクティブな対応力の強化が重要であると言える。

2.3 国や組織によるアクティブサイバーディフェンスに関する取り組み

アクティブサイバーディフェンスは、その概念の拡張性と含意の強さゆえに、国や組織ごとに異なる定義と運用方針が存在する。特に、「どこまでが防御か」「敵システムへの介入は許されるか」といった解釈を巡り、各主体の立場が分かれている。2.1 で述べたように、アクティブサイバーディフェンスには国際的に統一された定義がなく、国や組織によってその解釈と実践には大きな差異が見られる。これは、それぞれの法的枠組み、戦略的目標、技術的能力、そして脅威認識の違いを反映している。

● アメリカ合衆国

米国のアクティブサイバーディフェンスへの取り組みは、特に Department of Defense (DoD) とその関連機関において顕著である。Defense Advanced Research Projects Agency (DARPA) のアクティブサイバーディフェンスプログラムは、サイバーディフェンス担当者に「ホームフィールドアドバンテージ」を提供することに焦点を当てており、国防総省が管理するサイバー空間内で高度な敵対者

と直接的に関与し、脅威や脆弱性をリアルタイムで発見、定義、分析、軽減することを目指している。重要な点として、このプログラムは「本質的に防御のみを目的」としており、攻撃的な研究は除外されている。これは、定義された境界内での非常に積極的かつ技術的に高度な防御アプローチを示している。

一方、2018年に発表された Department of Defense (DoD) の”U.S. Department of Defense, Summary of the 2018 Department of Defense Cyber Strategy, September 2018”⁴において、サイバー空間を陸・海・空・宇宙に次ぐ「第五の戦闘領域」と定義し、積極的かつ継続的な防衛姿勢への転換を示している。特に重要な柱となるのが Defend Forward（前方防衛）と Persistent Engagement（持続的関与）の2つである。Defend Forward は、敵のサイバー活動が米国のシステムに到達する前に、それを起点で阻止するという前向きな防衛戦略であり、敵のネットワークや中継インフラに対して合法的かつ戦略的に介入することを含む。一方、Persistent Engagement は、敵と継続的に接触し、サイバー空間における優位性を維持することで、抑止と対応力を高める戦術である。これにより、DoD はリアルタイムでの脅威検知・対処能力を強化し、従来の受動的なセキュリティ手段では対処困難な高度な攻撃にも対応可能なようにしている。また、民間部門や同盟国との協力も重視されており、重要インフラ防護や国際的な安全保障枠組みの構築が進められている。さらに、すべての作戦において国際法や国内法を厳守し、正当性と透明性のあるサイバー行動が追求されている。これらを通じて、DoD は国家の安全保障を確保し、サイバー空間における優位性の持続を目指している。

● EU（欧州連合）

EU（欧州連合）は、2022年11月に発表した”EU Policy on Cyber Defence Communication”⁵において、サイバー空間における防衛能力の強化と加盟国間の協力促進を目的とした戦略的枠組みを示している。この政策は、「より強固な EU サイバー防衛のために共に行動する」「防衛エコシステムの確保」「サイバー防衛能力への投資」「共通の課題に対処するためのパートナーシップ」という4つの柱を中心に構築されている。

具体的な取り組みとして、EU Cyber Defence Coordination Centre (EUCDCC) の設立が挙げられる。EUCDCC は、防衛コミュニティ内での状況認識の向上とサイバー防衛活動の調整支援を担っている。また、軍事コンピュータ緊急対応チーム (milCERTs) 間の情報共有を促進するため、専用の通信ネットワークである MICNET の構築も進めている。さらに、EU 全体での検出、状況認識及び対応能力の強化を目的としたセキュリティオペレーションセンター (SOCs) のインフラ整備や、信頼できる民間サービス提供者を活用したサイバー予備軍の設立も推進されている。

このように、EU はアクティブサイバーディフェンスを含むサイバー防衛の強化に注力しており、防衛能力の向上及び加盟国間の連携強化、さらに国際的なパートナーシップの推進により、サイバー攻撃に対してより効果的な対応体制を整備している。

⁴ 「U.S. Department of Defense, Summary of the 2018 Department of Defense Cyber Strategy, September 2018」 <https://dodcio.defense.gov/Portals/0/Documents/Library/CyberStrategy2018.pdf>

⁵ 「EU Policy on Cyber Defence Communication」

https://www.eeas.europa.eu/sites/default/files/documents/Comm_cyber%20defence.pdf

- NATO（北大西洋条約機構）

NATO（北大西洋条約機構）は、2016年のNATOサミット（ワルシャワ）でサイバー空間を集団防衛の中心的な領域として位置づけ、アクティブサイバーディフェンスをその防衛戦略の重要な柱として採用している。また、2021年のNATOサミット（ブリュッセル）では、サイバーディフェンスの枠組みにおいてサイバー攻撃が状況に応じて集団自衛権（北大西洋条約第5条）の発動条件となり得ることを意味し、NATO加盟国は、サイバー攻撃が重大な影響を及ぼす場合、軍事的手段を含む全ての手段を用いて対応する可能性があることを示している。

また、NATOのサイバー防衛戦略は、技術的、戦略的、法的、運用的な4つの主要な側面から構成されており、これらの側面は、サイバー防衛能力の向上、国際法の適用、作戦の実行、そして教育・訓練の強化を目的としている。具体的には、NATOは「タリン・マニュアル」を策定し、サイバー空間における国際法の適用を明確にしている。

さらに、「ロックド・シールズ」などの実践的な演習を通じて、加盟国のサイバー防衛能力を強化している。これらの演習は、サイバー攻撃に対する実践的な対応能力を高めることを目的としており、AIや5G技術を取り入れるなど、最新の技術動向を反映している。

このように、NATOはアクティブサイバーディフェンスを通じて、サイバー空間における脅威に対抗し、加盟国の安全保障を強化している。今後も、技術の進展や新たな脅威に対応するため、戦略の見直しや能力の向上が求められると考えられる。

2.4 日本政府における「能動的サイバー防御」の考え方

日本政府では、サイバー攻撃の高度化・巧妙化に対応するため、「能動的サイバー防御」の実現に向けた検討が行われている。従来の防御手法では、攻撃の発生後に対処する受動的なアプローチが主流であったが、近年のサイバー攻撃の高度化や新たな攻撃手法に対しては、従来の方法では対応しきれないことが明らかとなっている。そこで、攻撃が発生する前、あるいは攻撃が拡大する前に、積極的に対策を講じる「能動的サイバー防御」という考え方が登場したのである。令和7年5月16日に成立した「サイバー対処能力強化法」⁶及び「同整備法」⁷は、国家安全保障戦略に基づく能動的サイバー防御の実現に向けた取り組みと考えられる。

能動的サイバー防御は、令和7年6月現在、衆議院で審議中の「サイバー安全保障を確保するための能動的サイバー防御等に係る態勢の整備の推進に関する法律案」⁸において、「外部からのサイバー攻撃について、これによる被害が発生する前の段階から、その兆候に係る情報その他の情報の収集を通じて探知し、その主体を特定するとともに、その排除のための措置を講ずることにより、国家及び国民の安全を損なうおそれのあるサイバー攻撃の発生並びにこれによる被害の発生及び拡大の防止を図ること」と定義さ

⁶ 「サイバー対処能力強化法」

https://www.cas.go.jp/jp/seisaku/cyber_anzen_hosyo_torikumi/pdf/houritsu.pdf

⁷ 「同整備法」 https://www.cas.go.jp/jp/seisaku/cyber_anzen_hosyo_torikumi/pdf/seibi_houritsu.pdf

⁸ 「サイバー安全保障を確保するための能動的サイバー防御等に係る態勢の整備の推進に関する法律案」

https://www.shugiin.go.jp/internet/itdb_gian.nsf/html/gian/honbun/houan/g21306007.htm

れている。これまでのサイバー防御の考え方と大きく異なるのは「被害が発生する前の段階から、その（攻撃の）主体を特定し、排除のための措置を講ずる」という点にある。

このような能動的サイバー防御の具体的な取り組みとして、以下の3つが挙げられる。

- 官民連携の強化

日本政府は、基幹インフラ事業者や重要施設がサイバー攻撃を受けた場合、政府と民間企業が連携して迅速に情報を共有し、協力して対応する体制の構築を目指している。これにより、攻撃者の活動を早期に把握し、必要な対策を迅速に講じることができるようにする。民間企業と政府が密接に連携することで、サイバー攻撃の影響を最小限に抑え、企業や国民の生活基盤を守ることが可能になる。

- 通信情報の利用

攻撃を未然に防ぐためには、通信事業者が持つ通信情報を活用することが有効であるとされている。日本政府は、通信事業者が提供する通信データを分析し、攻撃者が使用していると思われる悪質なサーバーやIPアドレスを検出し、攻撃を未然に防ぐための対策を講じる。これにより、攻撃者の活動が進行する前に、その兆候を捉え、対処することが可能となる。

- アクセス・無害化

最も進んだ能動的サイバー防御の手段として、攻撃者のサーバーに対する直接的な対応が検討されている。重大なサイバー攻撃に対しては、警察・自衛隊が攻撃者のサーバーに対し、その活動を無害化することを視野に入れている。具体的には、攻撃者が悪意のあるプログラムを仕込んでいるサーバーを特定し、そのプログラムを無効化、攻撃者のC2サーバーと通信を遮断するなどの対応が行われる。これにより、攻撃者の活動を封じ込めるとともに、攻撃の拡大を防ぐことができる。

これらの能動的サイバー防御の取り組みは、サイバー攻撃が発生した際に被害を最小限に抑えるための重要な手段である。しかし、その実施には法的な問題が伴う。例えば、現行憲法第21条に基づく「通信の秘密」や、不正アクセス禁止法など、サイバー防御活動の範囲に影響を及ぼす可能性のある法的枠組みが存在する。このため、能動的サイバー防御を実施するには、既存の法制度との整合性を確保し、必要に応じて法改正や新たな法律を制定することが求められている。

さらに、能動的サイバー防御を実施するための技術的な課題もある。例えば、攻撃を早期に検知するための高精度な脅威分析ツールや、攻撃者の活動を追跡・封じ込めるための高度な技術が必要となる。加えて、攻撃者の行動を正確に把握し、適切なタイミングで対応するための判断力と高度な専門知識も求められる。

日本政府は、能動的サイバー防御の実現によって、サイバー攻撃に対する防御力を強化し、国民生活や経済活動の安全・安心を確保することを目指している。この取り組みの実現により、サイバー攻撃に対する国の防御力は大きく向上し、社会全体のセキュリティレベルを高めることが期待されている。

第3章 脅威インテリジェンスの概要

本章では、脅威インテリジェンスの定義及びライフサイクルの概要を示し、その必要性や活用にあたっての前提事項について解説する。あわせて、脅威ハンティングの実施において脅威インテリジェンスが果たす役割や関係性についても述べる。

3.1 脅威インテリジェンスとは

3.1.1 脅威インテリジェンスの定義

「脅威インテリジェンス」という概念は、利用する組織や文脈により定義や運用方法が異なる。国内において「脅威インテリジェンス」という用語は、IoC (Indicator of Compromise) や攻撃者プロファイルなど、静的かつ断片的な脅威情報を指す意味合いで使われる場面が多い。実際、イングランド銀行 (Bank of England) が発行している文書⁹「CBEST Intelligence-Led Testing」において、“脅威インテリジェンスとは、標的とする組織の業務、ICT システム、またはそれらを経由して流れる情報に危害を加えたり、弱体化させたりすることを意図する悪意のある主体の身元、目標、動機、ツール、戦術に関する情報の収集と分析を戦略的に推進するプロセスの文脈に沿ったアウトプット”といった主旨のことが記載されており、脅威インテリジェンスの一部ではそのような脅威情報を指し示すケースも存在する。しかし、同文書において、“脅威インテリジェンスは周期的で進化的なプロセスであるため、次のインテリジェンス・タスクを計画する際には、適切な場合には、前回のサイクルの結果（成功と失敗）も考慮に入れる。”とも記載されており、循環的に脅威インテリジェンスを活用するというプロセスもまた脅威インテリジェンスであることを指し示している。

本レポートでは、「脅威インテリジェンスの教科書」¹⁰を参照し、以下のように脅威インテリジェンスを位置付ける。

“脅威（意図・機会・能力）に関する情報について、収集・加工・統合・評価・分析・解釈を行った成果物。あるいは当該成果物を作成するためのプロセス、組織を指すこともある。”

このように定義した脅威インテリジェンスは単なる情報の蓄積ではなく、戦術・運用レベルのセキュリティ意思決定において繰り返し活用される、サイクル型の支援基盤となる。したがって、その有効性は、データの質以上に、それをどう加工し、どの文脈に位置付けて、誰の意思決定を支援するかにかかっている。

3.1.2 脅威インテリジェンスの必要性

本レポートの冒頭でも述べた通り、近年、IT および OT システムの統合・高度化に伴い、サイバー攻撃の手法も複雑化・巧妙化している。加えて、攻撃が組織ごとの業種やシステム構成に合わせてカスタマイズされる傾向が強まり、汎用的な対策だけでは対応が困難な場面が増えている。

⁹ 「CBEST Intelligence-Led Testing」 <https://www.bankofengland.co.uk/-/media/boe/files/financial-stability/financial-sector-continuity/understanding-cyber-threat-intelligence-operations.pdf>

¹⁰ 「脅威インテリジェンスの教科書」/石川朝久 [著]/技術評論社

実際、欧州ネットワーク情報セキュリティ機関（ENISA）の報告や一般社団法人 JPCERT コーディネーションセンター（JPCERT/CC）の観測でも、ゼロデイ攻撃や多段階の標的型攻撃、正規ツールを悪用する Living-off-the-Land（LotL）型攻撃の増加が指摘されている。これらの攻撃は、従来のルールベースや署名ベースの検出手法では検知が難しく、画一的なフレームワークに依存する「コンプライアンス型アプローチ」では対応力に限界がある。この点について、「脅威インテリジェンスの教科書」では、「コンプライアンス型アプローチは、ベストプラクティスを網羅的に実施することには適しているものの、脅威環境の変化に追従しにくく、攻撃者の視点を踏まえた柔軟な対策立案には向いていない」と指摘している。

たとえば、NIST CSF 2.0¹¹でもリスクマネジメントにおける脅威環境の継続的把握と評価の重要性が強調されており、一律対応からの脱却が求められている。

このような背景から、自組織にとって本質的に重要なリスクとは何かを見極め、脅威を起点とした柔軟な対応方針の立案を可能にする「脅威ベース型アプローチ」が注目されている。本アプローチは、攻撃者の意図・能力・傾向といった外部の攻撃情報をもとに、自組織の脆弱性やビジネス構造と照らし合わせて、対策の優先順位や施策の適用範囲を文脈に応じて動的に調整する考え方である。「脅威インテリジェンスの教科書」ではこのアプローチについて、「攻撃者の視点に立ち、組織にとって真に重要な脅威を特定し、限られたリソースを適切に配分するために有効な戦略」であるとし、コンプライアンス主導の画一的な対策から脱却する手段として位置づけている。

このように、セキュリティ対策の設計においては、従来型のアプローチと脅威起点のアプローチの違いを理解し、適切に使い分けることが求められる。以下に、コンプライアンス型アプローチと脅威ベース型アプローチの特徴や利点・課題を整理した比較表を示す。

表 2 コンプライアンス型アプローチと脅威ベース型アプローチの比較

	コンプライアンス型アプローチ	脅威ベース型アプローチ
特徴	・ 法律や規制、業界標準などの外部基準への準拠を重視 ・ 評価基準はルールや要件の充足状況 ・ リスク評価も基準に沿ったものが中心	・ 組織特有の脅威を分析し、優先順位を設定 ・ 深刻度や発生確率に基づく対応 ・ 資産や業務への影響を重視したリスク評価
メリット	・ 外部基準に従うことで最低限のセキュリティを確保可能 ・ 導入が比較的容易で法的責任の履行に有効	・ 自組織に合った対策が可能で実効性が高い ・ 変化する脅威への柔軟な対応が可能 ・ 高度な防御策の導入がしやすい
デメリット	・ 組織固有のリスクに対応しきれない場合がある ・ 脅威の変化に対する対応が遅れる可能性 ・ リスクの深刻度にかかわらず均質な対策となり、資源配分が非効率となる可能性	・ 導入に専門知識が必要で運用が複雑化しやすい ・ 継続的な監視と分析が不可欠 ・ コンプライアンス上の説明責任を十分に担保しにくい場合がある

¹¹ NIST「Cybersecurity Framework 2.0」 <https://www.nist.gov/cyberframework>

こうした「脅威ベース型アプローチ」を実行性あるものとするうえで、中核的な役割を果たすのが、脅威インテリジェンスである。脅威インテリジェンスは、膨大かつ断片的な脅威情報の中から、自組織にとって意味のある脅威を特定し、分析・加工した上で、具体的な行動指針として知見を提供する。

さらに、脅威インテリジェンスは防御策の構築にとどまらず、プロアクティブなセキュリティ実践にも貢献する。たとえば、既知の攻撃パターンや攻撃対象の特徴を手がかりに、潜在する攻撃兆候を探索する「脅威ハンティング」においても、脅威インテリジェンスはその出発点として不可欠な役割を担っている（詳細は 3.3 にて後述）。

3.1.3 脅威インテリジェンスのプロセス

脅威インテリジェンスは、単なる情報の収集や提供にとどまらず、組織の意思決定を支援するために、計画的かつ体系的に運用されるプロセスである。

本レポートでは、IPA 産業サイバーセキュリティセンターから発行されている「脅威インテリジェンス導入・運用ガイドライン」¹²を参考に解説する。以下に示す図は、脅威インテリジェンスのライフサイクルを視覚的に表したものであり、各ステップが連続的かつ循環的に実施されることを示している。

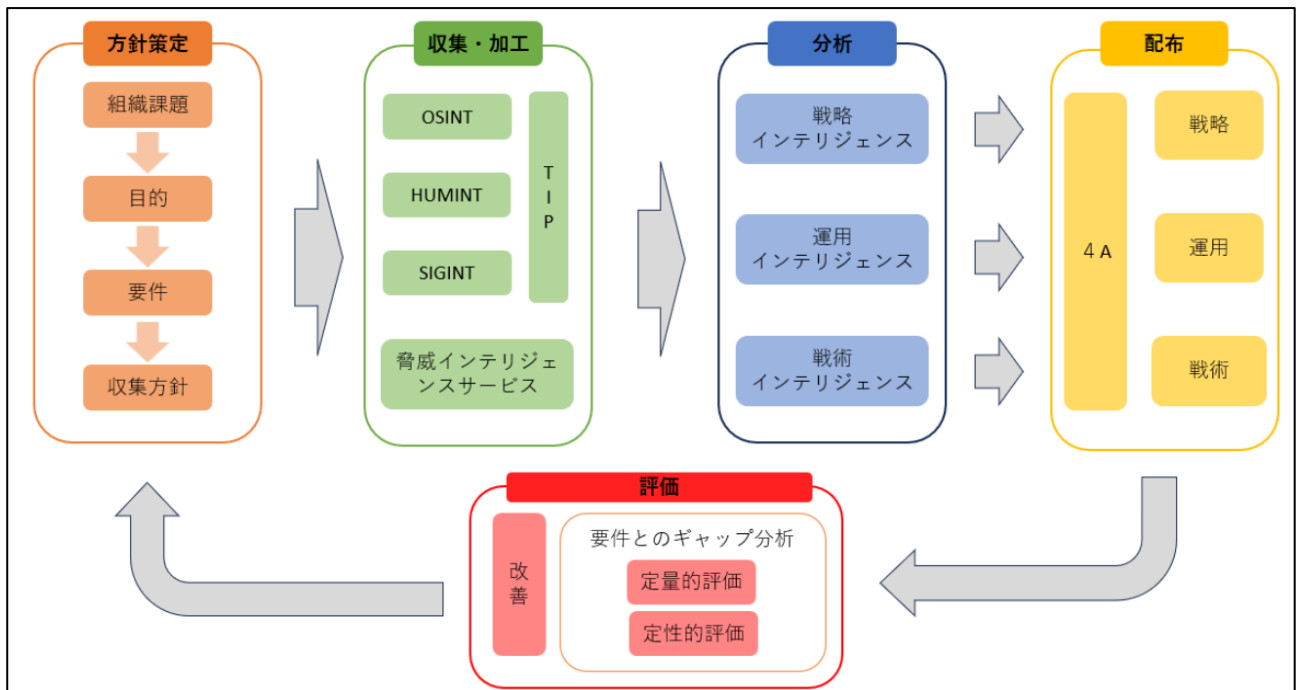


図 2 脅威インテリジェンスのライフサイクル
(出典：「脅威インテリジェンス導入・運用ガイドライン」)

このプロセスは一度きりの作業ではなく、継続的に見直しと改善を行うことが重要である。なお、図中の各ステップの詳細については、次節（3.2）にて個別に解説する。

¹² 「脅威インテリジェンス導入・運用ガイドライン」

https://www.ipa.go.jp/jinzai/ics/core_human_resource/final_project/2024/f55m8k0000003510-att/f55m8k000000358r.pdf

3.1.4 脅威インテリジェンスを利用する為の前提事項

脅威インテリジェンスを効果的に活用するには前提条件が必要となる。こちらについても様々な国際機関において定義されているが、概念や定義は類似しており、脅威インテリジェンスを意思決定に活かすためには、技術的基盤と組織的体制が整っていることを前提としている。

その中で、Robert M. Lee は「The Sliding Scale of Cyber Security」¹³において、サイバーセキュリティの取り組みを段階的に分類し、脅威インテリジェンスはその後半フェーズに位置付けている。すなわち、脅威インテリジェンスを導入するためには、それ以前の段階である Architecture（設計による防衛）および Passive Defense（自動化された防衛）の成熟であることを示している。

以下に、これら二つの要件について整理する。

① セキュリティを前提としたシステム設計と運用基盤の整備（Architecture）

システムやネットワークの構築・運用において、セキュリティを前提とした設計がなされていることが求められる。たとえば、以下のような取り組みが含まれる。

- ネットワーク分離やアクセス制御ポリシーの明確化
- 資産の可視化と管理
- 脆弱性管理やパッチ適用など、基本的なサイバー衛生の徹底

② 自動的かつ一貫性のある防御メカニズムの実装（Passive Defense）

人手を介さずに既知の攻撃に対処できるよう、一定の自動化された防御体制が機能していることが望ましい。具体的には、以下のような要素が該当する

- シグネチャベースの検知ルールやアクセス制御の自動適用
- 多層防御の構成
- EDR（Endpoint Detection and Response）や SIEM（Security Information and Event Management）による継続的な監視と対応

これらの基盤が成熟してこそ、脅威インテリジェンスの導入は単なる情報の収集にとどまらず、実践的な意思決定につながる知見の生成として真価を発揮する。

加えて、導入に際しては「どのような意思決定を支援するために脅威インテリジェンスを活用するのか」という目的を明確にし、文書化することが不可欠である。これは、脅威インテリジェンス・ライフサイクルにおける出発点、「方針策定フェーズ」に相当する。これらの目的を明文化し、組織内で合意形成を行うことによって、脅威インテリジェンス・プロセスは単なるデータ管理に終始せず、行動につながる知見として機能する仕組みへと昇華される。

3.2 脅威インテリジェンスのライフサイクル

3.2.1 方針策定

脅威インテリジェンスのライフサイクルにおける起点である「方針策定」フェーズは、インテリジェ

¹³ Robert M. Lee. 「The Sliding Scale of Cyber Security」 <https://sansorg.egnyte.com/dl/l2R3x1Ipxd>

ンス活動全体の方向性と目的を定める最も重要なプロセスである。この段階では、単に脅威情報を収集するのではなく、「どのような目的で、どのような知見を得たいのか」を明確にし、その後の収集・分析活動の軸を整えることが求められる。

「脅威インテリジェンス導入・運用ガイドライン」では、方針策定フェーズを以下の4つのステップで構成している。

① 課題抽出

自組織の課題を明確にする。守るべき資産や業務プロセスを把握し、自組織のセキュリティ体制や環境における課題を洗い出す。PESTLE 分析（政治/経済/社会/技術/法制度/環境）などのフレームワークを活用することも有効である。

② 脅威インテリジェンスの目的の設定

自組織の課題を起点として、「何のために脅威インテリジェンスを活用するのか」を明文化する。たとえば、「業界内の侵害情報をもとにプロアクティブに対処する」や「経営層への戦略判断材料として脅威動向を整理する」など、目的は組織の立場やニーズに応じて多様でよい。

③ インテリジェンス要件の策定

必要なインテリジェンス要件を明確にする設定した目的を達成するために、どのような種類の脅威情報（IoC、TTPs、脅威アクター情報など）を収集・分析すべきかを定める。「マルウェアの挙動」や「特定アクターの攻撃対象範囲」など、情報の具体性が精度を左右する。

④ インテリジェンス要件を満たす情報収集方法の検討

収集・分析手段の方針を定める必要な情報を得るために、どのような情報源（OSINT、HUMINT、パートナー連携など）を活用するか、また、どのようなツールや人材が必要かを含めて実行計画の原型を設計する。

このフェーズにおいて活動方針が曖昧なまま収集や分析に移行すると、組織にとって価値あるインテリジェンスを得られず、脅威インテリジェンス導入が単なる「情報の受け取り」で終わる形骸化を招くおそれがある。

このように、脅威インテリジェンスは“目的を持って活用する情報”であるという本質を踏まえ、最初の段階で目的と方針を明文化し、組織内で共有することが、成功の鍵となる。

3.2.2 収集・加工

脅威インテリジェンスにおける「収集・加工」フェーズは、前段の方針策定で定めた要件に基づき、さまざまな情報源から脅威関連データを収集し、分析可能な形式に変換・整備する工程である。

① 収集

情報の収集方法は主に以下の3種に大別される

- OSINT（Open Source Intelligence）

ニュースサイト、SNS、脆弱性情報共有サイト、フォーラム、技術系ブログなどから収集され

る公開情報であり、RSS フィードや Twitter API、検索エンジンの活用が代表例とされる。

- HUMINT (Human Intelligence)

ISAC などの情報共有コミュニティ、セキュリティベンダーからの通報、業界関係者とのヒアリングなど、人的ネットワークに基づく情報入手手段である。

- SIGINT (Signal Intelligence)

自社ネットワークに設置されたセキュリティ機器（ファイアウォール、EDR、WAF など）から出力されるログ、アラート、パケット情報などをもとにした技術的な取得手法である。

さらに、Threat Intelligence Platform (TIP) の導入により、複数の情報源を統合管理し、相関分析や自動整形を通じて情報の価値と再利用性を高めることができる。代表的な TIP には MISP¹⁴ や OpenCTI¹⁵ があり、IoC や TTPs のような構造化データを取り扱う。

また、外部の脅威インテリジェンスサービスの活用も増加している。これらは専任のアナリストが収集・加工・分析の一部を代行し、レポート形式で提供されるため、特に「ハッカーコミュニティ動向」「非公開キャンペーン情報」など、内部体制ではアクセス困難な情報に対する補完的手段として有効である。

② 加工

収集されたデータは、そのままでは分析に活用しづらいため、以下のような処理を施して整備される。

- フォーマットの標準化 (STIX/TAXII 形式への変換など)
- メタデータ付加 (信頼度スコア、検出時刻、検証ステータスなど)
- 重複排除・誤検知の除去
- 脅威アクターや攻撃手法との関連付け (エンリッチメント)

特に、TIP や SOAR (Security Orchestration, Automation and Response) といった自動化・連携ツールを活用した加工工程は、人的リソースの削減と品質向上の両面で効果が高く、今日の脅威インテリジェンス運用には不可欠な構成要素となっている。

このように、「収集・加工」フェーズは、インテリジェンスを分析可能な状態へと昇華させるための基盤的な工程であり、情報の信頼性や一貫性を担保するためには、綿密な設計と適切なツール選定が求められる。

3.2.3 分析

分析フェーズでは、収集・加工された脅威情報に対して組織固有の文脈を加え、意思決定に資するインテリジェンスへと昇華させる作業が行われる。このフェーズにおいて生成されるインテリジェンスは、活用対象や目的に応じて 以下の 3 つに分類される

- 戦略インテリジェンス (Strategic Intelligence)

経営層や CISO、セキュリティリーダーなどの上位の意思決定者を対象とし、中長期的なリスク評

¹⁴ MISP Project. <https://www.misp-project.org/>

¹⁵ OpenCTI <https://www.opencti.io/>

価やセキュリティ投資判断の支援に用いられる。APT グループの地理的傾向、業界横断的な攻撃トレンド、法制度の変化など、マクロな情報をもとに判断材料を提供する。

- 運用インテリジェンス (Operational Intelligence)

SOC マネージャー、セキュリティ管理者、IR チームなど現場に近い層を対象とし、短期から中期の防御態勢強化やシステム再設計に活用される。MITRE ATT&CK のような TTPs ベースの情報を利用し、攻撃者の傾向に応じた具体的対応に寄与する。

- 戦術インテリジェンス (Tactical Intelligence)

SOC オペレーターや CSIRT 担当者などの実務担当者を対象とし、日常的な運用における即時的な対応に用いられる。主に IoC (侵害の痕跡) や悪性 IP アドレス、ファイルハッシュ、URL など技術的な指標を中心とする。情報の有効期間は数時間から数日と短く、鮮度と信頼性が極めて重要となる。

以下は、これらの分類を目的/対象者/情報の粒度/寿命という観点から整理した表である

表 3 脅威インテリジェンスの分類

分類	対象者	目的	情報	寿命
戦略インテリジェンス	経営層、CISO、セキュリティリーダー	長期的なリスク評価、投資判断	APT の活動傾向、法整備の変化	長期
運用インテリジェンス	SOC マネージャー、IR チーム	中短期的な対策、検知精度向上	TTPs 分析、MITRE ATT&CK	中期
戦術インテリジェンス	SOC オペレーター、CSIRT 担当者	即時的な検知と防御強化	IoC、悪性 IP、URL、ファイルハッシュ	短期

このように、脅威インテリジェンスの分析フェーズでは、用途と活用者を明確に定めたうえで、適切な粒度・表現形式・流通経路で情報を加工・提供することが求められる。単なる分類にとどまらず、情報の性質と活用状況に応じた設計が、インテリジェンスの有効性を左右する重要な要素となる。

3.2.4 配布

配布フェーズでは、分析によって得られた脅威インテリジェンスを、実際の意思決定者や利用部門に届け、具体的な行動に結びつけることを目的とする。このプロセスが適切に設計されていなければ、たとえ高品質な分析がなされていても、組織としての有効な対応には至らない。

「脅威インテリジェンス導入・運用ガイドライン」では、配布品質の基準として「4A」と呼ばれる以下の 4 項目を評価軸として提示している

- Accurate (正確であること)

情報が裏付けのある事実に基づいており、誤情報や憶測が排除されている。

- Audience Focused (利用者目線であること)

受け手の技術レベルや意思決定階層に応じて、情報の粒度、形式、用語が最適化されている。

- Actionable (次のアクションにつながる)

情報を受けた際に「次に何をすべきか」が明確になっており、判断や対応に直結する内容となっている。

- Adequate Timing (適切なタイミングであること)

戦術情報においてはリアルタイム性、戦略情報においては定期性など、目的に応じた適切な提供タイミングが求められる。

これら 4A の観点は、分析から成果物を生成するプロセスと密接に関係しており、最終的に「誰に、何を、どのように届けるか」を調整する配布フェーズの役割が問われる。したがって、分析段階でこれらの要件に配慮しつつ、配布段階では形式/タイミング/媒体の最適化を図ることが理想とされる。さらに、配布内容はインテリジェンスの分類（戦略/運用/戦術）に応じた最適化が必要となる。

- 戦略インテリジェンス：経営層向けレポート形式、図表付き要約
- 運用インテリジェンス：ダッシュボード、スライドなどの視覚的整理
- 戦術インテリジェンス：IoC を STIX・TAXII 形式で提供、SIEM 連携に最適化

このように、配布手段も情報の性質に応じて切り替える必要がある。主な手段としては電子メールやダッシュボード、定期レポートのような形式もあれば、TIP（MISP、OpenCTI など）を用いたシステム連携、ISAC や JPCSRT/CC との外部情報共有・連携といった方法も存在する。

配布フェーズは単なる情報の送信にとどまらず、情報の「伝達」と「活用」の橋渡しを担う最終調整工程である。最終的な成果物が利用者の行動につながるよう、分析フェーズとの密接な連携と継続的なフィードバックを通じて、柔軟かつ文脈適合的な運用が求められる。

3.2.5 評価・改善

評価と改善フェーズは、脅威インテリジェンスのライフサイクルを継続的に改善し、目的に即した有効なサイクルとして進化させていくための重要なプロセスである。インテリジェンスは一度配布して終わりとするものではなく、その成果およびプロセス全体を見直し、次のサイクルの方針策定へとフィードバックを行うことによって、より実践的で精度の高い対応が可能となる。

本フェーズにおいては、まず配布したインテリジェンスが当初定義された要件（目的や想定利用者）を満たしていたかどうかを、定量的および定性的に評価する。「脅威インテリジェンス導入・運用ガイドライン」では、評価手法として再び「4A（Accurate / Audience Focused / Actionable / Adequate Timing）」の観点を適用することが推奨されている。これは単に配布物の品質を問うものではなく、インテリジェンスが実際の意思決定にどのように寄与したか、また実行された対策の効果や改善の余地などを総合的に確認するための評価基準である。

評価には、以下のような要素が含まれる

- 定性的評価

インテリジェンスの有用性に関する利用者からのフィードバック、実際の判断や行動への寄与内容、利便性や信頼性など。

- 定量的評価

適用されたインテリジェンスによって向上または改善された検知件数、誤検知や過検知率の変化、インシデント対応までの時間短縮など。

これらの評価結果およびギャップ分析を基に、ライフサイクルの各プロセス（収集、分析、配布）に改善の余地がないかを、再度 4A の観点から点検する。たとえば、収集対象が戦略インテリジェンスに偏重していた、分析手法が利用者の期待に合致していなかった、配布手段がタイミング的に適切でなかった、といった構造的な課題を抽出し、次の方針策定フェーズに反映することが重要である。

このように、評価と改善のフェーズは単なる事後レビューにとどまらず、各プロセスの有効性と整合性を見直し、脅威インテリジェンスの品質を維持・向上させるための重要な工程である。

3.3 脅威ハンティングとの関連性

脅威ハンティングと脅威インテリジェンスの関係性は、オランダの金融 ISAC が提唱した脅威ハンティングのフレームワークである TaHiTI (Targeted Hunting integrating Threat Intelligence) のホワイトペーパー、「Threat Hunting Methodology」¹⁶において、両者が密接に結びついていることが明確に示されている。特に、脅威ハンティングを効果的に行うためには、脅威インテリジェンスの知識が不可欠であるとされている。本ホワイトペーパーにおいて、両者の関係性を示す 3 つの主要な要素が挙げられている。

第一に、脅威インテリジェンスはハンティング活動の出発点となる。脅威アクターおよびその TTPs に関するインテリジェンスは、どのような脅威に優先的に対応すべきかを判断するための基盤となる。特に、標的となり得る自組織の業種や地域と密接な関連を有する脅威アクターに関する情報は、仮説を構築し、活動を開始する上で極めて重要な示唆を与えるものである。

第二に、脅威インテリジェンスはハンティング活動の文脈化および推進に寄与する。ハンティングの過程で観測された TTPs を脅威インテリジェンスと照合することにより、より広範な脅威の背景や関連する攻撃者像を明らかにすることが可能となる。このような相互作用は、仮説の再構築や調査対象の拡張につながり、結果としてハンティング活動の質的向上を促進する。

第三に、ハンティング活動は、新たな脅威インテリジェンスを生み出す重要な手段となる。実際のシステム環境で観測された未知の TTPs は、攻撃者の行動や能力を把握する上で貴重な情報源となる。こうした知見は、他組織や情報共有の枠組みを通じて共有され、他者の調査活動の出発点となることもある。結果として、複数の組織が連携し、より広範で深い攻撃者像の構築につながる。このような情報の共有と再利用の積み重ねが、全体としての防御能力の向上に寄与する。

このように、脅威インテリジェンスと脅威ハンティングは一方通行の関係ではなく、相互に補完・強化し合う動的な関係であるとされている。この関係性を正しく理解し活用することが、今日の高度化・巧妙化する脅威への対抗手段として極めて重要である。

¹⁶ TaHiTI 「Threat Hunting Methodology」

<https://www.betalvereniging.nl/wp-content/uploads/DEF-TaHiTI-Threat-Hunting-Methodology.pdf>

第4章 脅威ハンティングの概要

4.1 脅威ハンティングとは

本章では、脅威ハンティングの概念について、その定義、アプローチの種類及びライフサイクルについて述べる。

4.1.1 脅威ハンティングの定義

脅威ハンティングは、これまで多くの企業や組織によって独自の定義や方法論が提唱されてきた概念であり、その内容や手法には一定の幅が見られる。本レポートでは、脅威ハンティングの黎明期である2015年頃からその方法論を提唱し、脅威ハンティングフレームワークとして多くの業界関係者や専門家によって参照されている Sqrri 社の定義をもとに説明を行う。Sqrri 社は、2012年に米国マサチューセッツ州で設立されたサイバーセキュリティ企業（2018年、Amazon が買収）であり、特に脅威ハンティング分野で高い評価を受けている。Sqrri 社は、脅威ハンティングを”process of proactively and iteratively searching through networks to detect and isolate advanced threats that evade existing security solutions.”¹⁷（既存のセキュリティソリューションを回避する高度な脅威を検出・隔離するために、ネットワークを能動的かつ反復的に探索するプロセス）と定義しており、単なるアラート対応や定型的な分析にとどまらず、未知の脅威や巧妙な攻撃の兆候を自ら見つけ出す能動的な活動であることを強調している。

この定義の根底には、脅威ハンティングの中核的な考え方である「侵害前提」の視点が存在する。これは、「攻撃者はすでにネットワーク内に侵入している可能性がある」という前提に立ち、環境内に潜む痕跡や異常を能動的に検出しようとする姿勢を意味している。従来の受動的なセキュリティでは、インシデントが顕在化するまで攻撃を認識できないというリスクが存在するが、脅威ハンティングはその隙間を埋めるための戦略的手段として注目されている。

4.1.2 脅威ハンティングの目的と重要性

脅威ハンティングの目的は、すでに侵害されている可能性のあるシステム内から、攻撃者の痕跡をいち早く見つけ出し、被害の拡大を防ぐことにある。特に重要なのが、「Breach detection gap」という概念である。これは、攻撃者がシステムに侵入した時点から、その侵害が検知されるまでの期間を指す。前章でも紹介したオランダの金融 ISAC が提唱した脅威ハンティングフレームワークである TaHiTI のホワイトペーパー「Threat Hunting Methodology」でも、この侵害検出ギャップの重要性が強調されている。たとえば、Verizon が発行した2018年のデータ侵害調査報告書（DBIR）では、侵害の約68%が数か月間も検知されないままだったという。

脅威ハンティングは、このような検知の遅れを縮めることで、長期間の潜伏や大規模な被害につながる前に兆候を発見し、迅速な対応を可能とする取り組みである。

¹⁷ 「A Framework for Cyber Threat Hunting」 <https://www.threathunting.net/files/framework-for-threat-hunting-whitepaper.pdf>

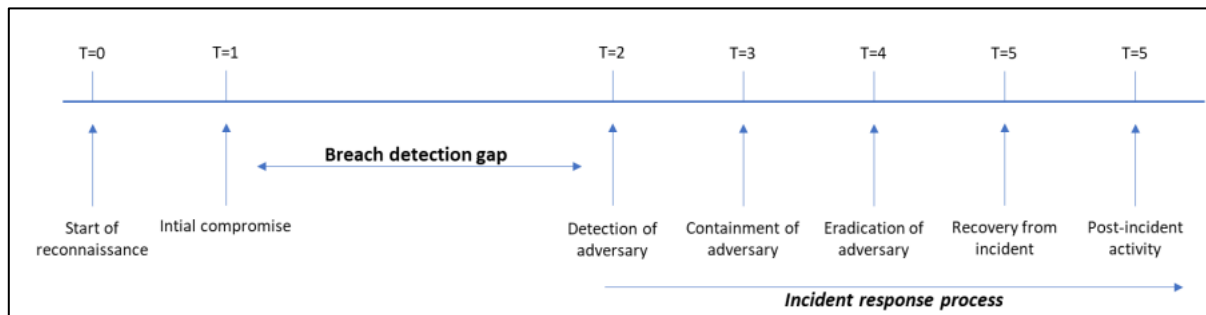


図 3 Breach detection gap

(出典：「Threat Hunting Methodology」)

このギャップが生まれる理由の一つは、攻撃者の TTPs が絶えず進化し、検知回避の巧妙さを増している点にある。特に高度な攻撃者は、検知を回避しながら長期的に組織内に潜伏する能力を持ち、既存のセキュリティツールでは発見が困難となる。このような「検知の盲点」に着目し、未知の TTPs や攻撃パターンを発見することも、脅威ハンティングが果たすべき重要な役割の一つと言える。

また、DX（デジタルトランスフォーメーション）の進展による IT 環境の複雑化も、脅威ハンティングの意義を高めている。クラウド、IoT といった多様な技術の導入は、従来のセキュリティ境界を曖昧にし、新たな攻撃経路を生み出している。これにより管理が行き届かない領域が拡大し、侵入者にとっての「狙い目」となっている。こうした領域を積極的にカバーする手段としても、脅威ハンティングはますます重要性を増している。

このような背景のもと、近年では脅威ハンティングをより戦略的かつ構造的に実施するためのフレームワークとして、MITRE Center for Threat-Informed Defense (CTID) が提唱した「Summitting the Pyramid (STP)¹⁸」が注目されている。STP は、従来の「Pyramid of Pain (痛みのピラミッド)」の概念を発展させたものであり、どのレベルのインジケータを起点とするかによって、検知の有効性や攻撃者への負荷が異なるかを定量的に示すものである。

とりわけ STP では、IP アドレスやハッシュ値といった低レベルの識別子よりも、攻撃の戦術や行動パターンといった TTPs レベルの指標に注目することが重視されている。これは、攻撃者が容易に変更できない振る舞いを起点とすることで、回避リスクを低減し、より効率的かつ本質的な検出が可能になるという脅威ハンティングのアプローチとも一致する。STP はその枠組みの中で、各種検知アナリティクスを「回避耐性 (evasion resistance)」と「誤検知率 (false positive rate)」の観点から評価できるため、ハンティング活動の成果を検知ルールの強化へとつなげることが可能となる。

優れた脅威ハンティングは、単に未知の脅威を発見するにとどまらず、既存の監視体制に内在する弱点を可視化する手段でもある。たとえば、通常の監視では見逃されがちなイベントを能動的に掘り起こすことやトリアージにおけるミスやアナリストの見落としによって検出されなかった事象を再検証することで、重要な示唆が得られる場合がある。加えて、そもそも脅威を検知・分析するために必要なログが収集されていないといった、監視設計そのものの不備を明らかにする重要な機会にもなる。これにより、ログ収集の設計見直しや監視ルールの改善といった具体的な対策が導かれ、結果として組織全体のセキュリティ成熟度の向上につながっていく。このように、脅威ハンティングは「未知の脅威への対

¹⁸ 「Summitting the Pyramid」 <https://ctid.mitre.org/projects/summitting-the-pyramid>

応」と「既存の監視体制の強化」という両面において、現代におけるサイバー防御の中核的な要素として、その重要性を一層高めていくと考えられる。

4.1.3 脅威ハンティングを実施する為の前提事項

効果的な脅威ハンティング体制を確立し、その能力を継続的に向上させるためには、①データ基盤の整備、②適切なツールの導入、③専門人材の育成・確保といった複数の要素が適切に整備されている必要がある。これらの要素は相互に関連しており、バランスの取れた発展が求められる。

また、自組織の取り組み状況を客観的に把握するためには、④ハンティング成熟度モデルの理解と活用も重要である。本節では、これら導入における前提条件などについて解説する。

① データ基盤の整備

脅威ハンティングの成果は、利用可能なデータの「質」と「量」に大きく依存する。脅威の兆候を分析によって見つけ出すためには、データの品質が極めて重要である。

具体的には、組織の環境全体を可能な限り網羅するため、多様なデータソースからの情報収集が不可欠となる。主要なデータソースとしては、EDR など詳細な挙動を記録するエンドポイントデータ、ネットワーク境界および内部の通信状況を把握するネットワークデータ（ファイアウォール、プロキシログなど）、各種セキュリティ製品が生成するアラート（IDS/IPS ログなど）、ユーザー認証やアクセス権限の状況を記録する IAM（アイデンティティアクセス管理）データ、クラウド環境における操作ログやリソースログであるクラウド環境ログ、さらには個別の業務アプリケーションの動作を記録したアプリケーションログなどが挙げられる。

加えて、外部から提供される脅威インテリジェンスフィード（既知の悪性 IP アドレス、ドメイン、マルウェアハッシュなど）や、組織内の資産情報（重要サーバーのリスト、管理者アカウント情報等）、ユーザーや部署に関するコンテキスト情報をこれらのログデータと組み合わせることで、検知された事象の重要度判断や分析の精度を大幅に向上させることが可能となる。

また、データ収集の際には、以下の観点が重要である。

- 正確性・完全性・適時性：不正確なログや遅延したデータは分析の信頼性を損なう。
- 正規化・標準化：異なる形式のデータを統一的に扱うことで、分析効率が向上する。

これらを実現するためには、データガバナンスと収集戦略の確立が必要と言える。

② 適切なツールの導入

高度な脅威を見つけ出すためには、膨大なデータを効果的に処理・分析できるツールの活用が重要である。主要なツールとして、以下のようなものが挙げられる。

- SIEM（Security Information and Event Management）
多様なログを集約・正規化し、相関分析を通じて脅威の兆候を検知・可視化するための分析基盤としてのツール。
- EDR（Endpoint Detection and Response）
PC やサーバーといったエンドポイントの挙動を常時監視し、プロセス実行や通信の挙動などから不審な活動の検知や封じ込めが可能な脅威検知ツール。
- NDR（Network Detection and Response in cybersecurity）
ネットワークトラフィックを監視・解析し、C2 通信やデータ流出といった不審な活動の特定

を目的とした脅威検知ツール。

- TIP (Threat Intelligence Platform)

複数のソースから脅威インテリジェンスを収集・整理し、正規化などを自動化することで、監視の効率と精度を向上させることを目的としたツール。

- UEBA (User and Entity Behavior Analytics)

ユーザーやエンティティ（サーバやルータなど）の通常行動をベースラインとして学習し、そこからの逸脱を検出することで内部不正や未知の脅威検出を目的とした脅威検知ツール。

効果的な脅威ハンティングには、膨大なデータを処理して分析し、脅威の兆候を可視化する適切なツールが不可欠である。これらのツールは、多様なデータを一元管理し、効率的な分析基盤を確保する必要がある。

これらのツールは、それぞれの役割に応じた強みを持っており、単体での利用に加え、連携させることで相互補完的な分析が可能となる点が重要である。多角的な視点からの分析により、より精緻な脅威の特定と対応が実現できると言える。

③ 専門人材の育成・確保

脅威ハンティングの成否は、データや分析ツールだけでなくそれらを駆使する人間の専門知識や経験にも依存する。

必要とされるスキルセットとして、以下のようなものが挙げられる。

- OS やネットワークの基礎に加え、攻撃手法への深い理解
- 大量データから仮説を立て、検証できる分析的思考能力
- インテリジェンスの文脈化と調査結果を効果的に伝達するコミュニケーション能力

また、こうした人材を単独で活動させるのではなく、複数の役割を持ったチームとして編成することで、専門知識を補完し合いながら柔軟な対応が可能となるようにすることが理想である。しかし、熟練した人材の確保は簡単ではなく、人的リソースが不足している場合は、外部ベンダーの支援やマネージドサービスの活用も現実的な選択肢と言える。

④ ハンティング成熟度モデル (HMM)

組織が自らの脅威ハンティング能力を客観的に評価し、計画的に強化していくためのフレームワークとして、David J. Bianco によって提唱された”Hunting Maturity Model” (HMM) ¹⁹が広く認知されている。このモデルは、組織が収集・活用しているデータの量と品質、脅威分析に関する能力、そしてハンティング活動における自動化レベルに基づいて、成熟度を5つの段階 (HMM0～HMM4) に分類している。

¹⁹ 「A Framework for Cyber Threat Hunting」 <https://www.threathunting.net/files/framework-for-threat-hunting-whitepaper.pdf>

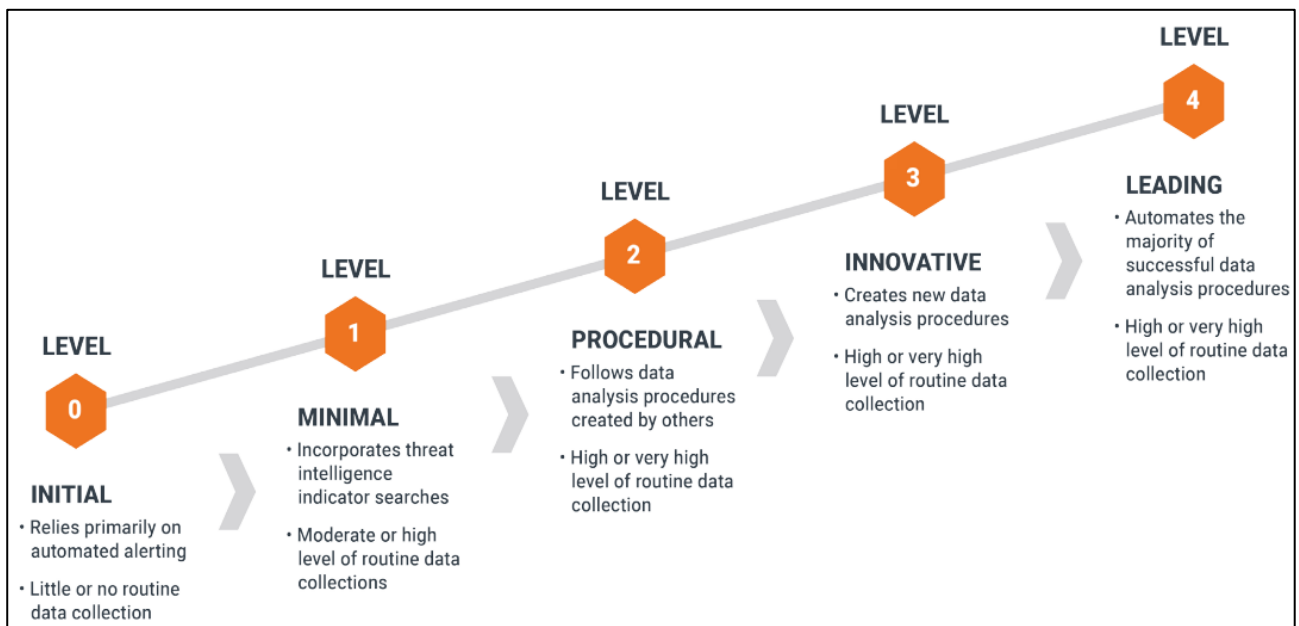


図 4 ハンティング成熟度モデル (HMM)

(出典：「A Framework for Cyber Threat Hunting」)

HMM の大きな特徴の一つは、評価軸として「どのようなツールを使っているか」ではなく、「どの程度の質と量のデータが日常的に収集され、それがハンティングに活用されているか」に重きを置いている点である。つまり、組織内で継続的に蓄積されるデータの範囲と深度が、成熟度レベルを決定づけるための要因の一つとしている。

ツールセットの有無や性能は補助的な要素にすぎず、質の高いデータと、それを解釈・活用できるアナリストのスキルがあれば、ツールの不足をある程度補うことも可能であるとされている。そのため、HMM においては、データの質とカバレッジに焦点が当てられており、分析に使用するツールについては明確な定義を設けていない。

また、ハンティングの成熟度が上がるにつれて、必要とされるデータはより詳細かつ文脈化されたものへと変化する。例えば、初期段階ではアラートやログといった限られた情報源に依存しているが、成熟が進むにつれて、構造化・非構造化を問わず多様なデータを統合的に活用し、既知・未知の脅威をより高度に分析・予測できるようになる。

このように、HMM は単なるスキル評価にとどまらず、組織がハンティング活動を持続可能かつ戦略的に拡張していくためのロードマップとして活用することができる。各レベルにおける要件を把握することで、自組織の現状を見極め、次の成熟段階に向けた具体的な改善方針を策定することが可能となる。

● HMM0 – Initial

組織は、主に IDS、SIEM、アンチウイルスなどの自動アラートツールに依存しており、脅威の検知は主にアラート対応に限られる。IT 環境からのデータ収集は自動アラートの作成に必要な最低限の情報以外に行われておらず、脅威ハンティングの専門知識があっても実施することは困難な状態にある。

● HMM1 –Minimal

組織は、依然として主に自動アラートに依存してインシデント対応を行っているが、ハンティング活動の土台となる IT データを定期的に収集している。このレベルでは、脅威インテリジェンスを活用した検知を志向しており、オープンソースやクローズドソースの脅威レポートを継続的に追跡している点が特徴と言える。新たな脅威が報告された際には、アナリストが関連する指標（IoC など）をもとに過去のデータを検索し、該当する痕跡がなかったかを確認することが可能である。このような履歴データの照会による検索が可能になることで、最小限ながらも脅威ハンティングが実施できるレベルと位置づけられている。

- HMM2 – Procedural

組織は、すでに確立されたハンティングフレームワークを取り入れ、それを日常的に実践しているレベルである。この段階の組織では、他者が作成した手順を学習・適用し、小さな修正を加えることはできるが、完全に独自の手順をゼロから設計する能力はまだない。

HMM2 の組織の特徴として、定期的な脅威ハンティングの実施が挙げられる（厳密なスケジュールに基づくものではない場合もある。）。また、最小頻度分析のような手法を活用するためには、多くのホストからのデータが必要なため、組織全体から大量のデータを日常的に収集しているのが一般的である。この HMM2 は、アクティブなハンティング活動を行う組織の中で最も広く見られるレベルであり、基本的な脅威ハンティングの体制が整っている状態とされている。

- HMM3 – Innovative

組織は、自律的かつ高度なハンティング能力を有していると言える。他者が開発した手順に依存せず、独自のハンティング手法を設計・文書化し、継続的に実行できる体制を備えている点が大きな特徴である。

このレベルの組織には、統計やログ分析に加え、可視化、機械学習などの多様な手法を理解し、適切に適用できるアナリストが少なくとも数名在籍していることが望ましい。重要なのは、これらのアナリストが複雑な分析技術を用いて、新たな脅威を発見するための手順を自ら構築し、それを再現可能なプロセスとして文書化している点にある。データの収集体制は HMM2 と同様に広範囲であるが、HMM3 では新たな分析手法の導入に伴い、収集対象のデータソースがより多様かつ高度になる傾向がある。時間の経過とともに、より高品質の分析に適したデータへと移行する姿勢が求められる。

このような組織は、実際の脅威アクターの行動を高い精度で発見し、迅速に対処できる可能性が高い。とはいえ、独自に開発したハンティングプロセスの数が増加するにつれ、それらを継続的に実行するためには人的リソースが必要となるため、スケーラビリティの課題に直面するリスクも存在する。したがって、能力の維持・拡張には、アナリストの増員や運用自動化の検討が必要である。

- HMM4 – Leading

組織は、HMM3 と同様に高度な分析能力と多様なデータソースを活用して脅威ハンティングを行っているが、両者の間には明確な違いがある。それは、成功したハンティング手法の自動化を実現している点である。

HMM4においては、アナリストが検出した脅威の兆候やプロセスを一過性の調査結果として終わらせるのではなく、それらを恒常的な自動検知ロジックとして組織の運用に組み込んでいく。これにより、アナリストは繰り返し手作業で同じ分析を行う必要がなくなり、新たな脅威への対応や、既存プロセスの改善に集中できる環境が整う。

このような継続的な自動化とプロセス改善により、HMM4の組織は敵対的な活動に対して極めて高い検出能力と対応力を発揮する。さらに、ハンティングを通じて得られた知見が検知プログラムに反映されるため、検知体制全体が絶えず進化していくことが期待される。

HMM4の状態に達した組織では、ハンティング活動が単なる探索行為に留まらず、検知インフラの高度化と組織全体のセキュリティ成熟度向上に貢献する持続的な戦略となっていると言える。

脅威ハンティング成熟度モデル（HMM）は、組織のハンティング能力を評価し、段階的な成長の道筋を示す指標である。特に、自動化の位置づけにおいて、HMM0とHMM4ではその意味が大きく異なるとされている。HMM0では、自動化されたアラートへの依存により受動的な対応にとどまるのに対し、HMM4では、成功したハンティング手法を運用化・自動化することで、継続的な改善と効率的な対応が実現可能となる。

また、HMMは単なる評価指標にとどまらず、ハンティングをこれから始めようとする組織にとってのスタート地点の設定や、すでに活動している組織にとっての改善計画の指針としても機能するとされている。明確な枠組みを持つことで、自らの立ち位置を把握し、次のステップに向けた取り組みを計画的に進めることが可能となる。HMMは、組織がより能動的で戦略的なハンティング体制を構築するための有効な指標なり得る。

4.2 脅威ハンティングの主なアプローチ

脅威ハンティングは、その起点や焦点の違いにより、複数のアプローチに分類される。本レポートでは、Sqrll社が提唱した5の分類を基に、それぞれのアプローチの特徴を明確化するとともに、TaHiTIが示す構造化・非構造化という視点を取り入れることで、各アプローチの位置づけや役割をより明確にできるよう独自に整理した。本節では、こうした観点から脅威ハンティングの主要なアプローチを体系的に解説する。

4.2.1 攻撃主導型（Attack based Hunting）

① 概要

攻撃主導型とは、攻撃者のTTPsに基づいて仮説を立て、その痕跡を組織内で探すアプローチである。MITRE ATT&CKフレームワークなどを参照し、特定の戦術（横展開や権限昇格など）に着目し、これらの行動が実際の環境で発生していないかを探索する。このハンティング手法は、攻撃者がどのように目的を達成しようとするかを予測し、その兆候を早期に発見することを目的としている。

② 特徴

攻撃主導型の主な特徴は以下の通りである。

- 攻撃者の視点と行動に着目
攻撃者がどのような TTPs を選択し、どのように目的を達成しようとするかを予測し、その痕跡を探す。
- MITRE ATT&CK フレームワークの積極的な活用
攻撃戦術/技術の分類、各 TTPs に関連するデータソースや検出方法の特定、ハンティングシナリオの作成、検知ルールの開発に至るまで、MITRE ATT&CK を共通言語および知識ベースとして広範囲に活用する。
- 攻撃の文脈理解
個々の TTPs を単独で評価するのではなく、それらが初期アクセス、永続化、権限昇格、内部偵察、情報窃取といったより全体像の中で、どのようにして利用されるのかという文脈で理解しようとする。

4.2.2 データ主導型 (Data based Hunting)

① 概要

データ主導型は、組織が収集/蓄積している膨大なセキュリティ関連データを統計的手法、高度な分析技術を用いて網羅的に分析し、その中から通常の活動パターン（ベースライン）から逸脱する異常な挙動（アノマリ）や、これまで認識されていなかった潜在的な脅威の兆候を能動的に発見しようとするアプローチである。この手法では、特定の既知の脅威情報や攻撃者の TTPs に必ずしも依存せず、データそのものが語る「通常とは何か」「異常とは何か」という観点から仮説を生成し、脅威を探索する。特に、UEBA は、このアプローチにおける代表的な技術として活用される。

② 特徴

データ主導型の主な特徴は以下の通りである。

- ベースラインからの逸脱検知
ユーザーアカウント、エンドポイント、ネットワークトラフィックなどの通常時の行動パターンを学習・確立し、そこから統計的に有意に逸脱する活動や、通常ではありえないパターンの出現を異常（アノマリ）として検知する。
- ML/AI 技術の応用
異常検知、クラスタリング、パターン認識などの ML/AI 技術を積極的に活用し、人間では発見が困難な複雑な相関関係や微細な兆候を捉えようとする。
- 未知の脅威発見への期待
既知のシグネチャやルールに依存しないため、これまで遭遇したことのない新しい攻撃手法や未知のマルウェア、内部不正の初期段階などを発見できる可能性がある。
- 仮説生成の起点
データ分析によって発見された異常なパターンや統計的な外れ値が場合によって、新たなハンティング仮説を立てる上での重要なトリガーとなる。

4.2.3 エンティティ主導型 (Entity based Hunting)

① 概要

エンティティ主導型は、組織内で特に保護すべき重要な「エンティティ」に焦点を当てて実施される脅威ハンティングのアプローチである。ここでの「エンティティ」とは、組織にとって最も価値の高い情報資産、特権アカウント、基幹業務システム、機密情報が保存されているサーバーやデータベース、あるいは組織の運営に不可欠な特定の重要な役割を担うシステムなどを指す。このアプローチでは、広大かつ複雑なネットワークの中で手当たり次第に脅威を探すのではなく、守るべき資産を明確に定め、その周辺で発生する異常な挙動や兆候に着目してハンティングを行う。調査範囲をあらかじめ絞り込むことにより、限られたリソースでも高い成果を上げやすい点が利点である。

② 特徴

エンティティ主導型の主な特徴は以下の通りである。

- 重要資産（エンティティ）中心のアプローチ
ハンティング活動の対象と優先順位が、組織にとって最も価値のある、あるいは最もリスクの高いエンティティに基づいて決定される。
- 深いコンテキスト理解の要求
対象となるエンティティの特性、通常の挙動、関連するビジネスプロセス、潜在的な攻撃経路などを深く理解することが、効果的なハンティングの前提となる。
- 範囲限定的かつ深掘り調査
ハンティングの対象が特定のエンティティに絞られるため、そのエンティティに関連するログやアクティビティに対して、より詳細かつ集中的な深堀調査を行いやすいという特徴がある。
- 内部脅威および標的型攻撃への有効性
特権アカウントの不正利用、内部関係者による情報窃取といった内部脅威、特定の重要エンティティを狙った標的型攻撃の検知に特に有効であり、UEBA 技術が活用されることがある。

4.2.4 インテリジェンス主導型（Intel based Hunting）

① 概要

インテリジェンス主導型は、IoC（侵害指標）や TTPs に関するレポートやフィード、脆弱性情報などの脅威インテリジェンスを出発点に、組織内に潜在する関連リスクを探索するアプローチである。攻撃者や手法に関する外部・内部の情報を分析し、特定の痕跡を優先的に調査することで、より戦略的かつ効率的なハンティングが可能となる。一方で、インテリジェンスの質や量に依存するため、信頼性を見極めと内部知見の蓄積が成功の鍵となる。

② 特徴

インテリジェンス主導型の主な特徴は以下の通りである。

- 脅威指標（IoC・IoA）に基づく調査
悪性 IP アドレス、ドメイン名、マルウェアのハッシュ値などの IoC（侵害指標）によって、既知の攻撃の痕跡を短時間で発見する「即効性」が期待できる一方で、攻撃者の行動パターンの文脈を示す IoA（攻撃指標）を活用することで、攻撃の兆候や挙動そのものにも着目したより深い分析と仮説構築が可能になる。
- インテリジェンス起点の探索

新たな脅威情報（最新の IoC、特定の攻撃グループに関する警告）の入手が、ハンティングを開始する「トリガー」となる。しかし、活動はそこで終わらず、その情報を基点として「この脅威は自組織にも存在するのか」「関連する他の攻撃の兆候はないか」といった仮説を立て、調査範囲を広げ、未知の被害や潜伏する脅威の発見を目指す。

- インテリジェンスの品質と結果

使用する脅威インテリジェンスの信頼性はハンティングの結果に直結する。低品質な IoC フィードは、誤検知や過検知を招くリスクが高く、分析の効率と正確性を損なう要因となる。そのため、信頼性の高い情報源の選定が重要と言える。

- 文脈理解による情報の有効性判断

得られた脅威インテリジェンスは、地政学的状況、業界特性、過去の攻撃傾向などを踏まえて、意味や関連性を判断することが重要である。人間による状況把握と意味付けが、仮説構築やハンティングの精度の向上につながる。

4.2.5 ハイブリッド型 (Hybrid Hunting)

① 概要

ハイブリッド型は、これまで述べてきた複数のアプローチを戦略的に組み合わせて実施する手法である。それぞれの長所を活かしながら、特定の脅威や資産に応じて柔軟に焦点を切り替えることが可能で、広範囲かつ深度のあるハンティングが実現できる。実際には、複数アプローチの併用が成功の鍵とされている。

② 特徴

ハイブリッド型の主な特徴は以下の通りである。

- 柔軟性と適応性

複数のアプローチを自在に組み合わせることで、精度、スピード、コストのバランスを取りながら、目的や状況に応じて最適化されたハンティングを実現可能。

- 網羅性の向上

複数の視点や手法を組み合わせることで多角的に捉え、検知率の向上を目指す。

- 相互補完性

一つの手法の限界を他の手法が補うことで、誤検出や見逃しのリスクを低減。

4.2.6 構造化・非構造化の視点によるアプローチ整理

脅威ハンティングは、その起点や進め方の違いにより、「構造化 (Structured)」と「非構造化 (Unstructured)」の2つのアプローチに分類される。この分類は、TaHiTI フレームワークにおいて採用されている考え方であり、ハンティング活動の明確な指針として活用されている。

表 4 脅威ハンティングの分類

分類	特徴	主な該当アプローチ
構造化ハンティング	事前に仮説や目的を設定して実施。調査範囲が明確で再現性が高い。	攻撃主導型 インテリジェンス主導型

	あらかじめ決められた枠組みに従って、計画的に調査を進める。	エンティティ主導型
非構造化ハンティング	広範なデータから異常を探索。自由度は高いが労力とコストが大きい。 あらかじめ決められた枠組みにとらわれず、人の直感を頼りにして調査を進める。	データ主導型
ハイブリッド型	構造化・非構造化の手法を柔軟に統合し、状況に応じて最適な方法を選択する。	上記の組み合わせ

構造化ハンティングとは、事前に立てた仮説に基づいて実施される脅威ハンティングである。仮説を出発点として、調査対象や範囲を明確に定め、系統的にハンティング活動を進める点が特徴である。具体的には、脅威インテリジェンスを活用した手法や、攻撃者の TTPs に着目する手法が該当する。また、組織にとって特に重要な重要資産などを起点に範囲を絞り込むエンティティ主導型のアプローチも構造化ハンティングの一種と位置付けられる。

一方で、非構造化ハンティングは、仮説から出発することもあるが、調査の進め方が定まっておらず、事前に構造化された手順に従わないために「非構造化」とされる。これは大量のログやシステムデータを分析し、異常値や未検知の挙動を見つけ出すことを目的とする。

TaHiTI フレームワークにおいて非構造化ハンティングは、調査の方向性があらかじめ定まっておらず、探索的であるため、多大な労力を要する割に、有益な成果を得るのが難しいという課題があるとされており、構造化ハンティングを開始するためのトリガーのひとつとして位置づけられている。

4.3 脅威ハンティングのライフサイクル

本節では、脅威ハンティングにおける一連の活動を体系的に捉えるための「ライフサイクル」について解説する。脅威ハンティングは多様な手法やアプローチに基づいて実践されており、組織やフレームワークによって構成要素や手順には差異がある。しかし、各種プロセスには一定の共通項が見られることから、本レポートでは2つの代表的なフレームワークを参照し、そこに共通する要素を抽出・整理したうえで、汎用的なライフサイクルモデルを提示する。

具体的には、脅威ハンティングのライフサイクルを体系的に理解するため、以下のフレームワークを参考とした。

- TaHiTI – Threat Hunting Methodology (Dutch Payments Association, 2018 年)
- PEAK – Prepare, Execute, Act with Knowledge (2023 年)

これらのフレームワークは、用語や分類に違いがあるものの、いずれも「計画」、「実行」、「完了」といった基本的な流れを共通して含んでいる。そこで本レポートでは、特定の用語や形式に依存せず、各フレームワークに共通するプロセスを軸にライフサイクルを再構成している。これにより、組織の規模や採用している手法を問わず、汎用的な視点から脅威ハンティングの全体像を把握しやすくすることを目的としている。こうした整理を踏まえ、以下に示すような汎用的な脅威ハンティングのライフサイク

ルモデルを提示する。本モデルは、実際のハンティング活動を段階的に捉えるための枠組みとして、今後の節で各フェーズの特徴や留意点を詳述する際の基盤となるものである。

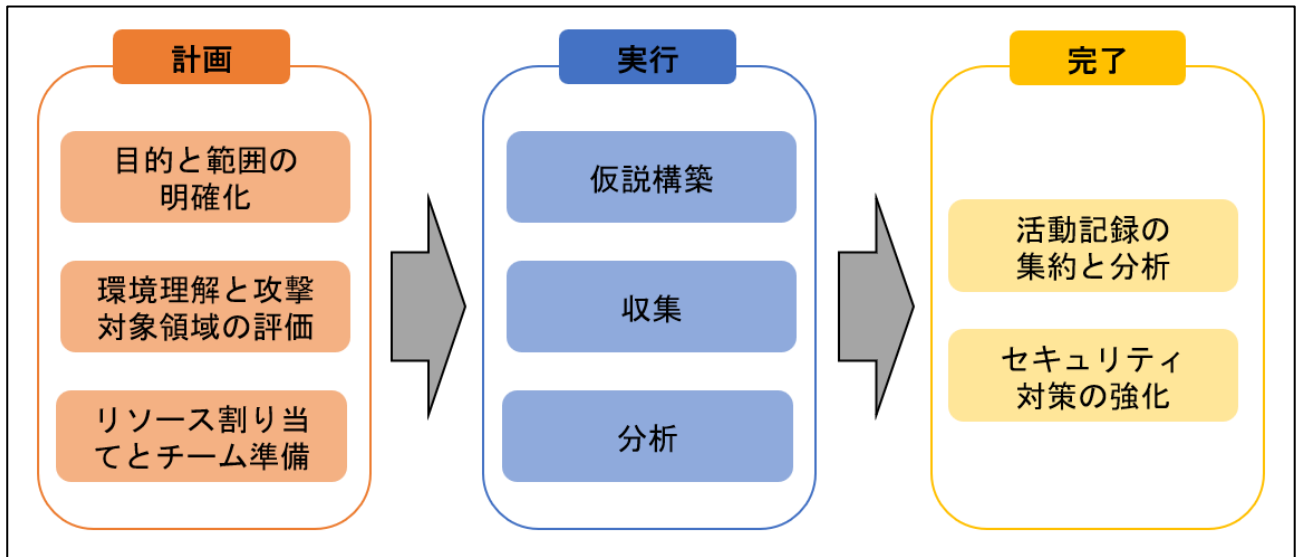


図 5 本レポートにおける脅威ハンティングのライフサイクル

以降では、このライフサイクルを構成する各フェーズ（計画、実行、完了）について順を追って解説していく。それぞれのフェーズで求められる視点や具体的な活動を明らかにすることで、脅威ハンティングの全体像を体系的に説明する。

4.3.1 計画

脅威ハンティングの最初の段階である「計画」は、活動の方向性を定め、リソースを効果的に活用するための基盤を築く重要なフェーズである。計画フェーズで考慮すべき主要な要素を以下に整理した。

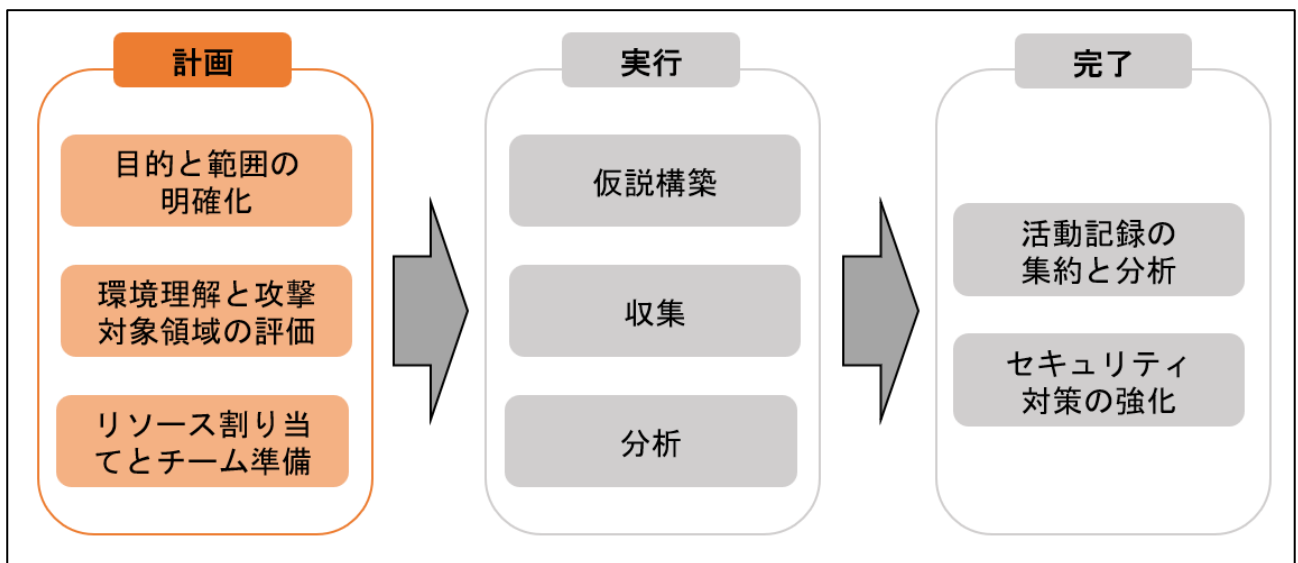


図 6 本レポートにおける脅威ハンティングのライフサイクル（計画）

① 目的と範囲の明確化

計画フェーズの出発点として、脅威ハンティングの「目的」と「範囲」を明確にすることが重要である。ここでの目的とは、どのようなセキュリティ上の懸念や脅威に対してハンティング活動を行うのかという方向性を定めることであり、組織が守りたい対象（重要資産など）や注目すべき攻撃手法（APT、ランサムウェアなど）を明確にする役割を果たす。目的を明確にすることで、限られたリソースの中で優先順位をつけ、効率的な活動が可能となる。

範囲の設定では、その目的に基づいて、調査対象を具体的に絞り込む。たとえば、特定のシステム、業務部門、ネットワークセグメント、あるいはユーザーグループなど、実施可能な単位で対象を限定する。このプロセスは、後に行う仮説構築やデータ収集の効率性と実現可能性を高めるためにも重要である。

このように、目的と範囲を明確にすることで、ハンティング活動が漠然とした探索にならず、現実的かつ戦略的なアプローチへとつながっていく。

② 環境理解と攻撃対象領域の評価

脅威ハンティングにおいて「異常」を見つけ出すには、まず「正常」を知る必要がある。そのため、組織が保有する IT 資産、通常の業務プロセス、通信パターンなどについて深く理解しておくことが重要である。こうした環境に関する理解が不十分な場合、誤検知や見逃しのリスクが高まり、活動の信頼性が損なわれ兼ねない。

加えて、すべてのシステムやネットワークを同じ深度で調査対象とすることは現実的ではない。そのため、リスクベースの観点から調査対象の優先順位をつけることも必要となる。外部との接点が多いシステム、ビジネスへの影響が大きいサービス、または既知の脆弱性を持つ資産などが、より重点的な調査対象となる傾向がある。

このような優先順位の設定にあたっては、現在観測されている脅威キャンペーンや既存の防御・検知体制、さらに脅威インテリジェンスに基づく TTPs の傾向なども考慮される。これらの考え方は、各ハンティングフレームワークでも、リスクベースの優先順位付けと攻撃対象領域の絞り込みが重要なステップとして位置づけられている。

③ リソース割り当てとチーム準備

ハンティング活動を計画する際には、設定した目的や範囲を踏まえて、仮説を検証するために必要となるリソースを整理しておくことが望ましい。ここでのリソースには、技術的なツール、対象データへのアクセス、チーム体制などが含まれる。

まず、仮説に基づいた調査を効果的に進めるために、適切なツールの選定が求められる。どのようなデータを扱い、どの観点から分析するかに応じて、利用するツールや手法を検討する。必要に応じてログ管理、ネットワーク監視、インテリジェンス活用など、目的に沿った機能を持つツールが選ばれることになる。

また、使用するデータへのアクセス権限についても、事前に確認・整理しておくことが望ましい。脅威ハンティングでは、システムログや通信履歴に加え、場合によっては個人情報や業務上の機密情報といったセンシティブなデータを扱うことがある。そのため、アクセスに関する内部規程や外部法令との整合性を確認し、必要に応じて法務部門などとの事前調整を行うことが重要である。適切な権限の付与と法的な確認を経たうえで、調査に必要なデータへのアクセスが確保されている状態を整えることで、後の工程を円滑に進めることができる。

さらに、ハンティングをチームで実施する場合は、作業分担や連携体制についても事前に調整しておくことが重要となる。仮説の立案、データの収集、分析、結果の記録など、各工程を明確に分担することで、全体の作業効率が向上する。

このような活動を通じて、仮説検証を円滑に進めるための体制が整い、実行フェーズに向けた基盤が構築される。

4.3.2 実行

実行フェーズは、前のフェーズで策定された計画に基づき、実際にデータを収集・分析し、仮説を検証して脅威の兆候を探索する、ハンティング活動の中核となる段階である。

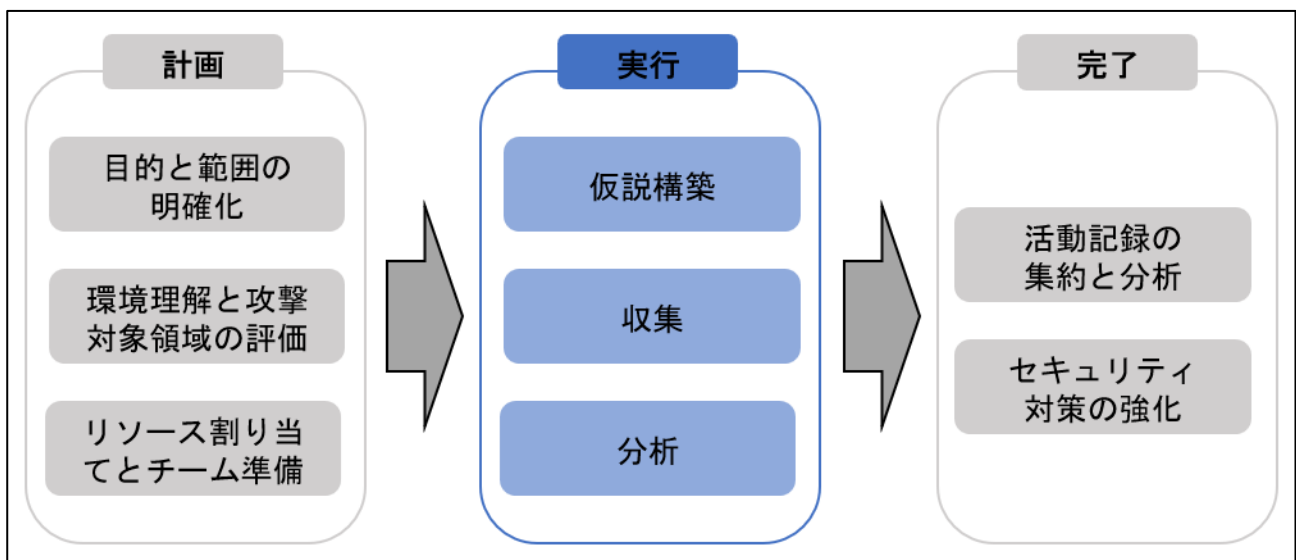


図 7 本レポートにおける脅威ハンティングのライフサイクル（実行）

① 仮説構築

脅威ハンティングにおいて重要な役割を担うのが「仮説構築」である。ここでは、自組織の環境内に潜んでいる可能性のある攻撃者の活動、特定の TTPs の利用、あるいは既に発生しているかもしれない侵害の兆候について、情報に基づいた推測や理論を構築する。効果的な仮説は、その後のデータ収集の範囲を絞り込み、分析作業の方向性を明確にすることで、ハンティング活動全体をより焦点を絞った効率的なものにする。本レポートでは、SANS のホワイトペーパーとして公開されている”Generating Hypotheses for Successful Threat Hunting”（Robert M. Lee および David J. Bianco 著）を参考に、有効な仮説構築のためのポイントを以下に解説する。

● 仮説構築の情報源

脅威ハンティングにおける仮説は、闇雲に立てられるものではなく、様々な情報源から得られる客観的なインテリジェンスや状況証拠に基づいて形成される。主要な情報源としては、以下のようものが挙げられる。

- 脅威インテリジェンス

最新の攻撃者の TTPs、サイバー攻撃キャンペーンの動向、新たに発見されたマルウェアの情報、既知の IoC（侵害の痕跡）や IoA（攻撃の痕跡）、特定の脅威アクターのプロファイルな

どは、現実世界の脅威に基づいた仮説を立てる上で重要である。

- 既知の脆弱性情報

自組織のシステムやソフトウェアに存在する既知の脆弱性に関する情報は、攻撃者がそれらをどのように悪用しうるかを予測する上で重要な手がかりとなる。

- 組織内部の情報とリスクプロファイル

自社のネットワーク構成、重要な資産の配置、通常の業務プロセス、過去のインシデント事例、そして組織固有のリスク評価の結果などを深く理解することは、組織特有の状況に合わせた仮説を立てるために重要である。

- 環境内の異常な活動

監視システム（SIEM など）から報告されるアラートや、過去のハンティング活動によって検出された確立済みのベースラインからの逸脱、原因不明のシステム動作なども、新たな仮説の起点として活用されることがある。

- 業種やシステムに関するドメイン知識

医療、製造、金融など、業界ごとに攻撃者が取る戦術や影響範囲は異なる。例えば、ICS 環境では一般的な IT 環境での軽微な挙動が重大なインシデントとなり得る。このような業界特有の知識を前提とした仮説も重要である。

これらの情報を組み合わせることで、脅威ハンターは「もし攻撃者がこの脆弱性を悪用したら、どのような痕跡がログに残るだろうか？」や「特定の脅威アクターが我々の業界を標的にしていると報じられているが、我々の環境ではどのような TTPs を使用して侵入を試みる可能性が高いか？」といった、具体的で検証可能な仮説を立てることが可能となる。

- 効果的な仮説の要素

生成された仮説がハンティング活動を効果的に導くためには、主に以下の要素が求められるとされている。

- 具体性

仮説は具体的であればあるほど、検証に必要なデータの収集や分析の焦点を明確に定めることができる。「何か悪いことが起きているかもしれない」といった漠然とした仮説では、調査範囲が曖昧になり、終わりの見えない探索に陥る恐れがある。これに対し、「特定のプロトコルを利用したデータ窃取が、特定のサーバーから、特定の時間帯に行われている可能性がある」といった、状況・場所・手法・時間などを明示した仮説は、効率的な調査を可能にする。具体性は、ハンティングの方向性を定め、活動を現実的かつ実行可能なものとするための重要な要素である。

- 検証可能性

仮説は収集可能なデータと利用可能な分析技術を用いて、その正否を検証できるものでなければならない。すなわち、理論として成立していても、それを裏付けるデータが取得できなければ、実際の脅威ハンティングには活かせないとされている。たとえば、「PowerShell 経由で C2 通信が行われている可能性がある」という仮説を立てた場合、検証のためには PowerShell の実行ログ、プロセスの親子関係、ネットワークトラフィックなどが必要となる。これは、センサーの設置状況、ログ収集の範囲、SIEM や EDR の機能といった、データ

収集基盤の整備状況にも依存する。したがって、仮説を立てる際には「どのようにして検証するか」「そもそも検証可能か」という視点を併せ持つ必要がある。検証できない仮説は、実行不可能な理論に留まってしまうため、ハンティング活動の実効性を著しく損なう恐れがある。

- 関連性

仮説は、組織の現在のリスクプロファイルや、実際に発生し得る脅威シナリオと高い関連性を持つものであることが重要とされる。たとえば、特定の国を拠点とする APT グループが、金融機関に対して「クラウドベースのファイル共有サービスを偽装した Spear Phishing」を用いて認証情報を窃取し、その後 VPN を経由して内部ネットワークに侵入しているという脅威インテリジェンスが得られたとする。この場合、単に自組織が金融業であるといった業種の一致だけでなく、VPN 接続の保護状況、メールセキュリティの制御、クラウドストレージの利用状況といった内部の構成要素を踏まえ、「同様の手法が自組織でも成立し得るかどうか」という視点で仮説を構築する必要がある。このような仮説は、脅威情報をそのまま適用するのではなく、自組織の技術的・環境的な文脈を考慮した上で、成立可能性を評価するものであり、より実効的な脅威ハンティングにつながる。

- 柔軟性と反復性

仮説は静的なものではなく、ハンティングの結果や新たな情報に応じて、継続的に見直されるべきものとされている。たとえ成果が得られなくても、仮説は将来のハンティングを支える知識資産となるため、文書化しておくことが推奨されている。実際、ハンターが調査の過程で複数の仮説を立て、検証可能で調査に適したものだけを追求するのは一般的な手法である。しかし、検証に至らなかった仮説についても、将来のハンティングに備えて文書化し、必要な技術や情報を調査しておくことが推奨されている。仮説の構築にはアジャイルな側面があり、反復と継続を前提としたプロセスである。たとえハンティングが失敗に終わったとしても、それにより環境内に脅威が存在しなかったことが確認できれば、セキュリティの強化に資する結果と考えられる。このように、仮説は一度の試みで成果を求めるものではなく、継続的に評価・更新されるべきプロセスとして位置づけられている。

● ハンターの直感と専門知識の役割

仮説の構築は、単に既知の情報を整理するだけでなく、脅威ハンター自身の経験や知見が重要な意味を持つとされている。経験豊富なハンターは、攻撃者の視点で状況を捉える力や、システムの微細な挙動を読み取る感覚を持っており、既存の手法では見つけにくい脅威の兆候を察知できる可能性がある。また、脅威インテリジェンスや各種ツールが提供する情報は、分析の出発点として有効ではあるが、人間の持つ創造性や批判的思考、また攻撃者の動機や行動原理に対する理解がそれを補完することで、未知の脅威に対する洞察につながる場合もある。ただし、専門知識や経験が先入観を生むこともあるため、仮説構築においては客観性を意識し、新しい視点や異なる仮説にも柔軟に対応できる姿勢が求められる。

② 収集

脅威ハンティングにおける「収集」は、前の「仮説構築」段階で立案された仮説を検証するために不可欠なデータを特定し、収集して利用可能な状態にするプロセスである。この段階の成功は、収集

されるデータの質、網羅性、そしてそれらのデータへ効率的にアクセスし分析できるかどうかに大きく左右される。つまるところ、適切なデータがなければ、どんなに優れた仮説も検証することができない。

- 仮説に基づくデータソースの特定

脅威ハンティングにおいて、まずは、仮説を検証するための適切なデータソースを特定する必要がある。明確で具体的な仮説が構築されていれば、その仮説に基づき、必要となる情報の種類や観察すべき対象が自ずと明確になり、データソースの選定は比較的容易になる。

データソースを特定する際には、まずそのデータにアクセス可能であるかを確認する必要がある。たとえば、ある種のログが有効化されていなかったり、収集対象のシステムがモニタリングの対象外であったりする場合、仮説の検証自体が不可能となる。

次に、データの形式や粒度が仮説の検証に適しているかを評価する必要がある。たとえば、イベントログがサマリ形式でしか記録されていない場合や、ネットワークトラフィックがヘッダ情報のみにとどまりペイロードが取得されていない場合には、調査に必要な詳細な挙動を把握することが困難となる。また、ログの保存期間が短く、仮説が対象とする過去の挙動が既に削除されている場合も、検証に支障をきたす。

さらに、仮説が複数の視点からの情報によって成立するケースでは、それぞれの視点に対応したデータソースを網羅的に洗い出すことが重要である。たとえば、「PowerShell を経由した C2 通信が行われた可能性がある」といった仮説では、PowerShell の実行ログだけでなく、当該プロセスの親子関係、ユーザーの操作履歴、通信先 IP アドレスのトラフィックログなど、複数の情報を突き合わせる必要がある。

このように、仮説に対して「どのデータを使って、どう検証するか」という視点を持ち、実際の環境で利用可能なデータソースを適切に選定することが、脅威ハンティングの実効性を大きく左右すると言える。

- データ収集と品質確保

脅威ハンティングを有効に実施するためには、必要なデータを集めるだけでなく、それを分析可能な状態で整備し、高品質に保つことが重要とされている。ここでは、仮説に基づいた調査を支えるためのデータ収集と品質管理の基本的な観点について述べる。

- データ収集

データ収集においては、仮説に関連する情報が適切な期間・粒度・形式で利用可能であることが前提となる。特に、ハンティングではリアルタイムまたは過去の時系列データを参照する場面が多く、適切な期間にわたって必要な粒度でデータが保存されているかが重要な検討事項となる。

また、近年の IT 環境はオンプレミスとクラウドの混在やデバイスの多様化により、データが分散して存在し、収集や統合にコストがかかるという「データグラビティ」の課題も指摘されている。したがって、無理なくデータを集約・活用できるよう、事前にデータの流れや保管場所を考慮した収集設計を行うことが推奨される。

- データの集約と整備

脅威ハンティングの効率化には、異なるデータソースを一元的に扱える基盤の整備が必要となってくる。以下のような処理が推奨されている。

- 集約：複数のソースから収集したデータは、SIEM などの中央プラットフォームに集約

することで、横断的な相関分析が可能となる。

- 正規化：形式や構造が異なるログを共通のフォーマットに変換することで、異なるソリューション間で得られたデータの整合性を確保し、分析を容易にする。たとえば、2つの異なる EDR から取得したログに基づいて統合することが挙げられる。
 - エンリッチメント：IP アドレスに対する地理情報やレピュテーション情報、ユーザーに紐づく部署情報など、付加価値のある情報を追加することで、分析に必要な文脈が強化される。
- データ品質の管理

収集したデータの品質管理は、ハンティングの成果に直結する要素である。具体的には、以下の観点からの確認が求められる。

- 完全性：ログに欠損がなく、すべての必要な項目が揃っているか。
- 正確性：データの内容が正しく記録されているか。
- 時刻の整合性：異なるシステム間でのタイムスタンプが正確かつ同期されているか。

また、手で収集されたデータや一時的に外部から取り寄せたログについては、形式の不統一や不整合が頻繁に発生する。そのため、以下のような事前処理が必要となる場合がある。

- JSON 形式のログを CSV に変換する
- タイムスタンプを UTC に統一する
- 無効なレコードや欠損値を除去する

このように、整備された状態でデータを扱うことが、仮説の検証における分析の信頼性を高め、誤検知や調査の行き詰まりを防ぐうえで重要であると言える。脅威ハンティングにおいては、データの量だけでなく、質と整備の程度が成功の鍵を握っている。

ハンティング中にデータソースの不足や品質の問題が明らかになった場合は、それをフィードバックし、データ収集戦略を継続的に見直し、改善していくことが重要である。

③ 分析

脅威ハンティングにおいて、分析フェーズは収集したデータを精査し、立てた仮説を検証するための重要なプロセスである。この段階では、直感と論理的思考、そして適切な分析手法の活用が求められる。目的は、仮説を裏付ける証拠、あるいはそれを否定する材料をデータの中から見つけ出すことにある。

分析には、さまざまな手法がある。例えば、最小／最大発生頻度の分析では、特定のイベントの異常な出現傾向を検出することができる。クラスタリングは、類似した動作を行うエンティティをグループ化し、通常とは異なる振る舞いを明らかにするのに有効である。また、可視化を用いることで、時間的・空間的なパターンや通信の偏りなどを直感的に把握しやすくなる。これらの手法は状況に応じて使い分けられ、仮説の精度を高めるために反復的に適用されることが多い。

また、分析の結果として導かれる仮説の検証には、以下の3つの結末が考えられる。

1 つ目は仮説が証明される場合で、セキュリティインシデントの存在が明らかとなる。これにより、対応措置や新たな検出ルールの作成が可能となる。

2 つ目は仮説が反証される場合で、実際には問題が存在しなかったことが確認される。この結果も、検知のカバレッジ不足の発見や新たな改善点の抽出といった副次的な価値をもたらす。

3つ目は結論に至らない場合である。これは、データの不足、分析手法の不適切さ、あるいは仮説自体の不明瞭さなどが原因で、判断が困難なケースである。このような場合、調査を一旦立ち戻り、仮説や分析範囲の再定義を行ってハンティングを繰り返す必要がある。

このように、分析と仮説検証の工程は単なる一度きりの確認作業ではなく、柔軟に方向を変えながら深掘りしていく反復的なプロセスである。的確な分析と判断により、脅威の実態に迫ることができ、脅威ハンティングの目的である早期発見と予防に大きく貢献すると言える。

4.3.3 完了

脅威ハンティングの「完了」フェーズは、実施された一連の活動から得られた、知見や教訓を組織全体のセキュリティ体制の継続的な改善へと結びつけるプロセスである。このフェーズは、単に個別の活動を終結させるだけでなく、将来の脅威に対する予防、検知、対応能力を進化させ、脅威ハンティングプログラムそのものの成熟度を高めていくためのフィードバックループとしての役割を担う。

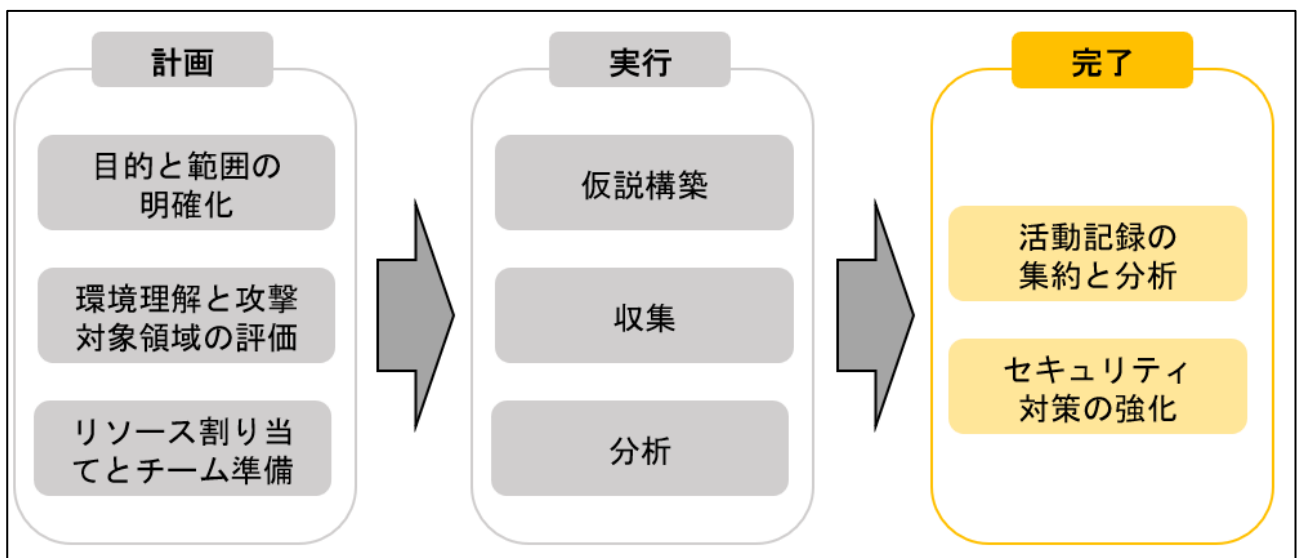


図 8 本レポートにおける脅威ハンティングのライフサイクル（完了）

① 活動記録の集約と分析

完了フェーズは、活動全体の記録とその振り返り・分析が中心的な役割を果たす。ハンティングの結果を詳細なレポートとして文書化し、どのような仮説を立て、どのようなデータを用いてどのような分析を行ったか、またどのような結論が導き出されたかを明確に記録する。

文書化には以下のような内容が含まれる。

- 立案された仮説とその根拠
- 分析対象としたデータソース、使用ツール
- 具体的な分析手順とプロセス
- 脅威が確認された場合の対応措置とその結果
- 発見された脆弱性や設定ミス
- 予防策の改善やログ記録に関する推奨事項
- ハンティング活動中に直面した課題や、プロセスの有効性に関する評価

これらの情報は、客観的な事実に基づいて整理・評価され、ハンティング活動の成功要因や課題、分析手法の有効性、全体の効率性などを振り返るために活用される。また、これらの記録はナ

レッジベースとして蓄積・共有することが重要であり、他のハンターや関係チームが今後の調査で参照できるように整備しておくことが望ましい。

また、分析結果は、「脅威ハンティングレポート」として関連するステークホルダー（技術チーム、SOC・IR チーム、セキュリティ管理者、経営層など）に報告される。読者の立場に応じて、技術的な詳細とビジネス的なインパクトの両面をバランスよく記述し、「Need-to-Know の原則」（知る必要のあるべき人に伝える）に基づいて適切に共有する必要がある。

② セキュリティ対策の強化

完了フェーズのもう一つの重要な目的は、ハンティングの成果を活用してセキュリティ対策を具体的に強化することである。これは、単に分析を終えるのではなく、得られた洞察を組織の防御力向上に直接つなげる実践的なステップとなる。

● 検出能力の向上

ハンティングで発見された IoC（Indicator of Compromise）や IoA（Indicator of Attack）、攻撃者の TTPs をもとに、SIEM や EDR などの検知ルールを新たに作成・強化する。これにより、将来的に同様の脅威が発生しても、より早期に、かつ自動的に検知できる体制を構築する。

● 監視体制と可視性の改善

活動中に判明したログの死角や監視対象外の領域、ログ品質の問題（欠損、不整合など）に対しては、ログソースの追加、収集設定の最適化、正規化プロセスの見直しなど、技術的改善を計画・実施する。これにより、次回以降のハンティングやインシデント対応の精度が大幅に向上する。

● インシデント対応体制の見直し

ハンティング中に新たな攻撃手法や脅威パターンに直面した場合は、それに対応できるよう、既存のプレイブックや対応フローを更新し、実践的な対応力を高める。

● 戦略とプロセスの再評価

仮説の立て方、ツール選定、分析アプローチの有効性を振り返り、次回のハンティング戦略の改善に役立てる。プロセス全体の洗練と標準化は、脅威ハンティングの成熟度を高める鍵となる。

このように、「完了」フェーズでは調査結果を記録・分析するだけでなく、それをもとに具体的なセキュリティ改善策を講じることが重要である。このプロセスを繰り返し継続することで、脅威ハンティングは単発のイベントではなく、組織全体の防御力を底上げするサイクルへと昇華していく。

本章では脅威ハンティングの基本理論を整理してきたが、それだけをもって組織の活動として実践へと昇華させるのは容易ではない。そこで第 5 章では、より柔軟で複合的な視点を取り入れたアプローチを模擬した攻撃シナリオに対してどう機能するかを検証する。この検証を通じて、運用効率や導入の現実性を踏まえ、どのような手法が実践的かつ有効となりうるのかを考察する。

第5章 脅威ハンティング実践検証

5.1 検証設計

5.1.1 検証の目的

本検証は、実際の組織への導入を見据え、運用効率などを考慮したアプローチを検証し、より実践的かつ現実的な脅威ハンティングの在り方を明らかにすることを目的としている。従来、脅威ハンティングは、専門的な知識や人的リソースを要する高度な活動とされており、そうした高度な手法をそのまま実際の組織運用に適用するには、さまざまな課題が伴うと考えられる。そこで本検証では、効率的な運用を可能にする手法、特に再利用性の高いプロセスや検出ロジックの有用性を確認する。また、限られた監視体制下でも有効に機能するかどうかを含め、現場への適用可能性を検証し、組織にとって導入が現実的なアプローチを考察している。

5.1.2 検証における仮説と前提

本検証で対象とするアプローチは、「攻撃主導型」と「ハイブリッド型」の2つである。この前提のもと、本章では以下の仮説に基づいて検証を行う。なお、ここでの「仮説」は、検証の設計上用いる前提としてのものであり、脅威ハンティングそのものの方法論としての「仮説」とは異なる観点であることに注意してもらいたい。

攻撃主導型は、MITRE ATT&CK フレームワークなどに基づき、既知の攻撃者の TTPs をベースにして脅威を探索する手法である。このアプローチは、攻撃の意図やパターンを明確に想定した上でログを追跡するため、既知の攻撃手法への即応性や検出ロジックの再現性や運用の標準化に優れている。一方で、未知の攻撃や環境固有の異常には対応しにくいという制約がある。

これに対し、ハイブリッド型は、インテリジェンス主導型の枠組みを土台に、データ主導型（異常な通信や挙動を統計的に検知）、エンティティ主導型（システムや資産の重要度や環境特性を考慮）、および攻撃主導型（攻撃者が使用する TTPs）といった複数のアプローチを統合することで、より柔軟かつ多角的な脅威検知を目指すものである。このアプローチは特に、既知・未知を問わず潜在的なリスクに対応しやすく、実環境において高い実用性が期待される。

本検証においてこの2つのアプローチを選定した背景には、現実の組織運用におけるバランスの重要性がある。攻撃主導型は運用の簡素化や効率化に有効である一方、ハイブリッド型はより高度な検知を可能にし、未知の脅威や高度標的型攻撃への対応力を補完する役割を担う。これらの手法を検証することで、組織の規模やリソースに応じた最適な導入アプローチを明らかにすることが可能になると考え、検証対象として採用している。

5.1.3 検証環境

検証は仮想基盤上に構築された企業ネットワークを模した閉域環境にて実施する。この環境には、Active Directory を中心とした Windows ドメイン、複数の業務端末、ファイルサーバー、ログ収集基盤が含まれる。各端末には EDR を導入し、Sysmon による詳細な挙動監視も併用している。通信監視にはネットワークパケットキャプチャシステムを設置し、ホスト・ネットワークの双方から分析を行える体制を整えた。ログ分析は SIEM を中心に実施し、検出ロジックの実装・検証も同一環境で完結可能とした。監視体制としては、実際の SOC 運用を模しつつ、ログ分析などによる脅威ハンティングが実施可能

な環境を構築した。

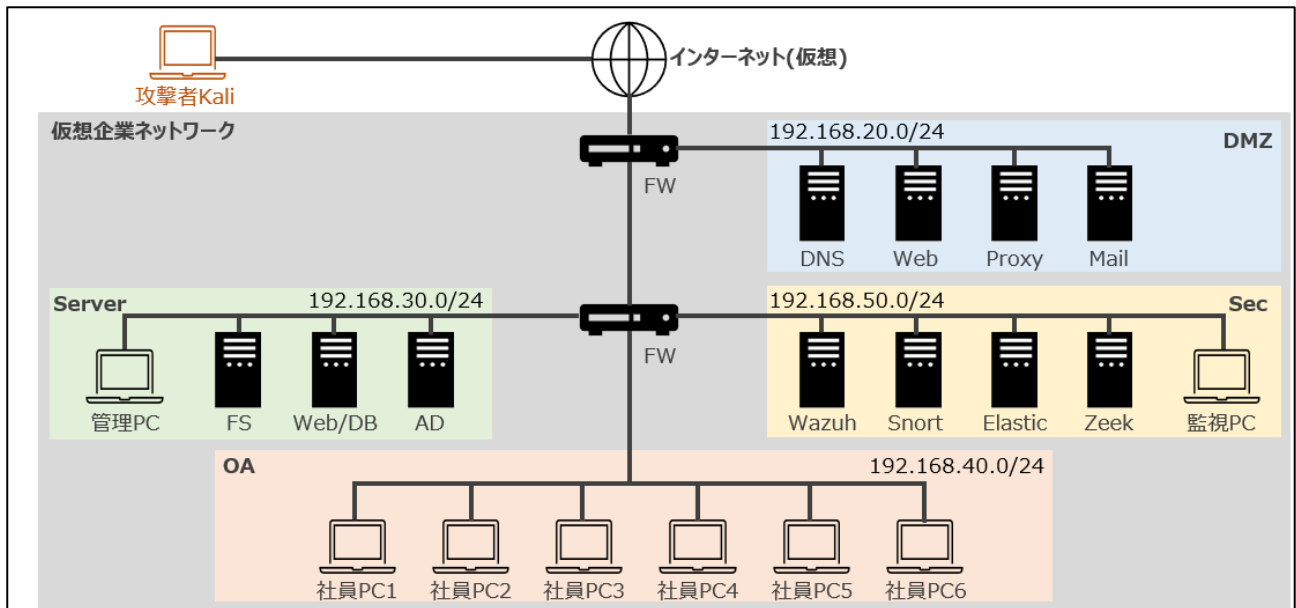


図 9 検証環境

5.2 攻撃主導型の脅威ハンティング検証

5.2.1 脅威シナリオ

本節で用いる脅威シナリオは、MITRE ATT&CK Evaluations (ATT&CK Evaluations: Wizard Spider & Sandworm) にて公開されているエミュレーションプランを参照して構築した。対象となる攻撃グループ「Wizard Spider」は、主に金銭的動機を背景としたランサムウェア攻撃を展開しており、TrickBot や Ryuk などのマルウェアを活用していることで知られる。このシナリオは、初期アクセスから権限昇格、水平展開、C2、データの収集と暗号化に至るまで、複数の TTPs を段階的に再現する構成となっている。

シナリオの起点は、「被害者となるユーザーが悪意のある文書ファイルを開封すること」による初期感染である。この文書は、外部からマクロコードをロードし、MSHTA 経由でコマンドが実行される仕組みとなっており、T1203 (Exploitation for Client Execution) や T1133 (External Remote Services) に該当する手法が使用されている。続いて、TrickBot などのマルウェアに感染させることでバックドアが確立され、攻撃者は環境内に持続的にアクセスが可能となる。この段階では、T1059 (Command and Scripting Interpreter)、T1027 (Obfuscated Files or Information) などの技術が用いられている。

また、特筆すべき点として、攻撃者は AD 認証情報の取得やサービスの無効化、データの暗号化に加え、セキュリティ製品の回避行動にも注力しており、これらは T1003 (Credential Dumping)、T1562 (Impair Defenses) などに対応している。さらに、データの暗号化後には T1486 (Data Encrypted for Impact) が適用されており、最終目的である身代金要求へと至る構成である。

このように本シナリオは、実際の攻撃チェーンを包括的に再現しており、多様な技術的手法が含まれているため、攻撃主導型のアプローチを検討するうえで有用である。

5.2.2 実践プロセス

本検証における攻撃主導型のプロセスは、MITRE ATT&CK の中でも、実際の攻撃において頻繁に使

用されるテクニックに着目することで、効率的かつ現実的なアプローチを実現することを目指した。その出発点として参照したのが、MITRE Engenuity および Center for Threat-Informed Defense (CTID) が「Sightings Ecosystem」²⁰プロジェクトの一環として公開している「Top 15 Techniques」である。

この「Top 15 Techniques」は、世界中のセキュリティパートナーから収集された 600 万件以上の TTPs の観測データに基づき、実際のサイバー攻撃において最も頻繁に確認された 15 の ATT&CK テクニックを特定したものである。すべての ATT&CK テクニックを網羅的に検出対象とすることは非現実的であることから、この 15 のテクニックは、実際の脅威に基づいた優先順位付けの合理的な出発点となると考えた。本検証ではこの知見を活用し、限られたリソースの中で有効な脅威ハンティング手法の確立を目指している。

表 5 MITRE Top15 Techniques

No.	Technique Name	Technique Name (日本語)	Technique ID
1	Command and Scripting Interpreter	コマンドおよびスクリプトインタプリタの使用	T1059
2	Obfuscated Files or Information	ファイルまたは情報の難読化	T1027
3	Ingress Tool Transfer	侵入ツールの転送	T1105
4	Modify Registry	レジストリの変更	T1112
5	Indicator Removal	インジケータの削除	T1070
6	User Execution	ユーザー実行	T1204
7	Hide Artifacts	アーティファクトを隠す	T1564
8	Process Injection	プロセスインジェクション	T1055
9	OS Credential Dumping	OS 資格情報のダンプ	T1003
10	Remote Services	リモートサービス	T1021
11	Data Encrypted for Impact	影響度に応じて暗号化されたデータ	T1486
12	Replication Through Removeable Media	リムーバブルメディアを介したレプリケーション	T1091
13	System Information Discovery	システム情報の検出	T1082
14	Windows Management Instrumentation	WMI (Windows 管理インストルメンテーション) の利用	T1047

²⁰ 「Sightings Ecosystem」 https://center-for-threat-informed-defense.github.io/sightings_ecosystem/

15	Impair Defenses	防御機能を弱める	T1562
----	-----------------	----------	-------

本検証では、各テクニックに対して複数の検出ロジックを用意し、多角的な視点から脅威の検出を試みている。以下で、検出ロジックの一部について説明する。

- Remote Services (T1021) の場合

T1021 は攻撃者がリモートサービスを利用して横展開や侵入後の操作を行うテクニックであり、検証環境を踏まえて以下のサービスを焦点とした。

- T1021.001 – RDP (Remote Desktop Protocol)
- T1021.002 – SMB/Windows Admin Shares
- T1021.003 – Windows Remote Management (WinRM)

上記以外にも wmic の利用状況を監視対象に含めるほか、Windows のログオンイベントにおける logon_type:3 (ネットワークログオン) も注視することで、リモート操作の兆候を効果的に検知できる。こうした多様なデータソースとイベントタイプを組み合わせる検出クエリを設計し、攻撃活動の痕跡を拾い上げることを狙いとしている。

また、Sysmon の導入有無にかかわらず検出可能なよう、Windows の標準イベントログ (セキュリティイベント ID 4624、4648 など) も活用し、ログ監査ポリシーの適切な設定を前提としている。加えて、通常の業務で利用される正当なリモート接続と攻撃の違いを判別するため、ユーザーアカウントの種類や接続タイミング、使用されたコマンドの特異性などの文脈情報を詳細に分析している。

- Command and Scripting Interpreter (T1059) の場合

T1059 は、攻撃者がスクリプトやコマンドラインインターフェースを活用して、初期アクセス後の操作や持続性の確保、さらにはペイロードの実行など多様な目的で利用するテクニックである。本検証環境では、以下の代表的なサブテクニックを焦点としている。

- T1059.001 – PowerShell
- T1059.003 – Windows Command Shell (cmd.exe)
- T1059.007 – JavaScript / JScript

これらの手法は、いずれも正規ユーザーによる運用にも使われるため、単純な実行検知のみでは攻撃との判別が困難である。そのため、以下のようなイベントおよび検出指標を組み合わせる検出アプローチを採用している。

本検証では、疑わしいスクリプト実行の検出に向けて、Sysmon のイベント ID 1 (プロセス作成) や、Windows のセキュリティログにおけるイベント ID 4688 を活用している。これにより、PowerShell や cmd.exe といったスクリプト実行環境の起動状況を分析した。特に、PowerShell においては、-enc (エンコード済みスクリプトの実行) や -nop (プロファイルの無効化)、Invoke-Expression (任意コードの評価・実行) といったオプションの使用を検出の指標とした。

このように、本検証の攻撃主導型は、観測された頻度の高い攻撃技術を優先しつつ、多面的なログ分析と検出ロジックを駆使することで、現場の運用負荷を抑えながら高い検知率を追求するアプローチを

実践しようとしている。

5.2.3 検証結果

本節では、攻撃主導型による脅威ハンティングを実施した結果について述べる。検証の結果、検証環境および再現した攻撃チェーンにおいて、特定の挙動を捉えることは一部可能であったものの、攻撃主導型単体の運用では、十分な精度で攻撃的な活動を検出することは困難であった。特に、環境固有の業務フローや通常のシステム挙動を考慮しないまま汎用的な検出クエリを適用するだけでは、多くの攻撃的活動を見逃す結果となった。

たとえば、「Command and Scripting Interpreter (T1059)」においては、PowerShell や cmd.exe の実行ログをベースに検出を試みたが、コマンドライン引数に直接的な悪意の兆候が含まれていない場合では、単一のログからの判定が非常に困難だった。

そこで、SIEM の統計機能を活用し、プロセスの親子関係において発生頻度が極端に低い組み合わせに注目するアプローチを採用した。結果として、mshta.exe や wsmprovhost.exe を親とする PowerShell や cmd プロセスの出現が、本検証の初期侵入や横展開と一致する兆候として確認されたが、これはあくまで限定的なシナリオに依存した成功例に留まり、実際の運用環境で機能するかの判定は難しい。

「Remote Services (T1021)」に関しても同様で、wsmprovhost.exe を用いた WinRM 通信のうち、社員ネットワークから特定の管理サーバーへの管理者ユーザーでのアクセスなど、通常想定されない通信パターンに着目することで、攻撃的な挙動の一部を検出できた。

本検証では、攻撃主導型のアプローチに限定して検証することを想定しており、その他のアプローチを活用することは検討していなかった。一方で、事前に用意した汎用的な検出クエリを機械的に適用しただけでは、攻撃チェーンに含まれる多くの挙動を検出するには不十分であった。特に、通信先情報やプロセスの実行環境、業務の流れといった組織固有の文脈を加味しない抽象的なクエリでは、正常なログによるノイズが多く含まれ、脅威の痕跡を効果的に抽出することが困難であった。

次節では、検出ロジックの限界や改善点、さらに実運用を見据えたハンティング手法の検討について考察を行う。

5.2.4 検証結果に対する考察

今回の検証を通じて、事前に想定された攻撃手法（いわゆる TTPs）に沿った汎用的な検出ロジックを適用するだけでは、脅威を捉えることは難しいということが明らかになった。たとえば、正規の業務でも頻繁に使われる PowerShell や cmd.exe の実行は、それ自体を検出対象とした場合、正常なログが多数含まれるため、ノイズが多く脅威の検出にはつながりにくい。このため、親プロセスや実行タイミングなど、プロセスが発生した背景（文脈）に着目した追加分析が必要であることが示された。

この「文脈」を捉えるためには、攻撃者の手口に着目する攻撃主導型のアプローチに加え、別のアプローチを組み合わせる必要がある。具体的には、ログの前後関係や、「どの端末で」「どのユーザーが」といった情報から重要資産を守る視点を持つエンティティ主導型、さらに、大量のログデータから「普段とは違う動き」を統計的に見つけ出すデータ主導型の視点である。これらのアプローチを有効に機能させるためには、組織における「正常な通信・挙動とは何か」を明確に把握しておくことが前提となる。各アプローチを組み合わせることにより、単体では問題ないと見なされる挙動も、特定の状況下における「異常」として検知できるようになると考えられる。

本検証では当初、検出アプローチをそれぞれ個別に適用しても一定の成果が出ると想定していたが、

実際にはそれぞれが補完し合う関係にあり、単体ではカバーしきれない部分があることが分かった。効果的な脅威ハンティングを行うには、各アプローチの強みを活かしながら柔軟に組み合わせることが求められる。複数の視点を統合した「ハイブリッド型」の分析手法こそが、現実の脅威に対応できる実践的な手段であると考えられる。

5.3 ハイブリッド型の脅威ハンティング検証

5.3.1 脅威シナリオ

本検証では、MITRE ATT&CK Evaluations にて公開されている FIN7 グループのエミュレーションプランをベースとした脅威シナリオを用いた。FIN7 は、主に金融機関や小売、ホスピタリティ分野を標的とする高度な攻撃者グループであり、金銭的動機に基づいた攻撃活動を長年にわたり展開している。本シナリオは、エンドユーザーのメール操作をトリガーとした初期アクセスから始まり、スクリプトベースのバックドアの展開、ラテラルムーブメント、情報収集に至るまで、段階的に攻撃行為を再現する構成となっている。

シナリオは、標的ユーザーがメール経由で送付された悪意ある添付ファイル（LNK ファイル）を開封することで開始される。これにより、PowerShell を通じて初期の C2 チャネルが確立され、T1059.001（PowerShell）や T1204.002（Malicious File）といった技術が使用される。その後、攻撃者は Windows Management Instrumentation（WMI）や Windows Credential Editor（WCE）などを用いて環境内での特権昇格や情報収集を実施し、T1547（Boot or Logon Autostart Execution）、T1003（Credential Dumping）などの技法を適用している。

また、FIN7 は商用ツールとオープンソースツールを組み合わせた高度な技術を用いており、本シナリオでも C2 通信におけるステルス性の確保や、検知回避のためのコマンド難読化、ログ消去などが含まれている。加えて、T1071.001（Web Protocols）や T1560.001（Archive via Utility）を通じたデータ圧縮・送信も試みられており、組織内部での持続的活動と外部流出のリスクが示される構成となっている。

このシナリオの特徴は、エンドポイント上での動作が一見すると正規操作に近く、検出が困難である点にある。そのため、本検証では、これら一連の行動に対して検出ロジックがどのように適用できるかを確認するとともに、攻撃の初動から終盤にかけての各フェーズをどこまで追跡・可視化できるかを検証することを一つの目的としている。これらを踏まえ、ハイブリッド型のアプローチの技術的・運用的な有効性を検討する材料として本シナリオは有用であると考えた。

ただし、脅威ハンティングは本来、「脅威を見つけ出す」という部分が強調されるべきであり、本検証で実施しているような攻撃全体の流れを把握する取り組みは、脅威ハンティングの主眼ではないことを補足しておく。

5.3.2 実践プロセス

本検証におけるハイブリッド型は、インテリジェンス主導型のアプローチを土台に、TTPs、アノマリ分析、資産の重要度や環境特性を組み合わせることにより、より現実的で適応力のあるアプローチを目指して実施した。本手法は、多層のかつ動的な視点を持ち、実際の運用において柔軟に適用可能なアプローチとして構成されている。まず初めに、ハンティングの出発点としてベースとなる脅威インテリジェンスを選定した。今回の検証では、FIN7 に関する脅威レポートや公開された IoC、行動パターンを

もとに、攻撃グループの特徴を明らかにした。その後、このインテリジェンスを MITRE ATT&CK フレームワークにマッピングし、類似する TTPs（特にサブテクニク）をリストアップ。これにより、攻撃者が利用する可能性の高い行動パターンを予測可能な形で整理した。

次に、実際のハンティングの焦点となる「メインスコープ」を選定した。今回は「初期侵入（Initial Access）」、「認証アクセス（Credential Access）」、「横展開（Lateral Movement）」、「収集（Collection）」及び「流出（Exfiltration）」の5つを主要な関心領域とし、それらに関連する TTPs や検出指標を掘り下げていくこととした。特にこれらは、攻撃の進行度や被害範囲を大きく左右するフェーズであり、早期の検知が重要である。

メインスコープの設定後、それに関連する異常挙動（アノマリ）や IoC に基づき、攻撃の可能性に関する仮説を構築した。たとえば、横展開に関しては、通常見られない SMB 通信、異常な時間帯におけるログオン、WinRM や PsExec などのリモート実行が仮説の出発点となった。仮説ごとに複数の観測ポイントを定義し、複合的な検知ロジックに落とし込むことで、誤検知を抑えつつ検出精度を高める工夫を行っている。

さらに、ハンティングのリソースを最適化するため、組織内の資産情報を考慮した。具体的には、システムの重要度（業務への影響度、機密性、外部接続の有無など）をもとに優先度付けを実施し、同一仮説内でも分析対象を絞り込むことによって、現場の負荷軽減を図った。この資産価値の視点は、限られた人員や時間で最大限の効果を発揮するための現実的な判断軸となる。

このように、本アプローチは、多層的な視点を融合させることで、攻撃の兆候を把握し、効果的に対応するための高度な実践プロセスを提供している。

5.3.3 検証結果

本検証において、ハイブリッド型を実施した結果、攻撃の各段階における攻撃痕跡を高い精度で検出できることが確認できた。検出結果は攻撃的活動を示すログが主であり、ノイズは少なく、分析に要する時間も短縮されたことから、実務での迅速な対応を可能にする有効な手法である可能性が示された。また、重要資産に焦点を当てた分析によってリソースの効率的活用が実現され、現実的な運用面での負荷軽減にもつながった。

検出結果の具体例として、認証情報取得を示す lsass.exe や NTDS、SAM ファイルへの不正アクセスに加え、Windows 正規管理ツールの一つである PsExec を使用した横展開の痕跡が明確に検出された。これには、異常なリモートプロセス実行のログや不自然なネットワーク接続が含まれ、SMB、WinRM を含む複数の横展開手法の監視と併せて効果的に検出された。

ただし、本検証はあくまで模擬環境における限定的な結果であり、実際の企業環境における多様な条件や変動要因を完全に再現したものではない。特に、Windows 正規管理ツールの使用などは攻撃者にとって検知回避の有効な手段であり、複雑かつ多様な環境下での検出はより困難なことが予想とされる。

5.3.4 検証結果に対する考察

本検証結果から、ハイブリッド型のアプローチは、複数のアプローチを組み合わせることで攻撃の多様な段階に対して効果的な検出を実現できることが示された。特に、誤検出を抑制しつつ、高い真陽性率を維持できた点は、本手法の大きな強みであると考えられる。

また、重要資産に優先順位をつけた分析は、限られた人的リソースを効率的に運用するうえで有効であることが分かった。一方で、本アプローチは適用に際して高度な知識と経験を要求するため、属人性

の高さが課題として浮き彫りになった。アナリストのスキルに依存する部分をいかに標準化し、ナレッジの共有や自動化によって補うかが今後の重要なテーマである。

さらに、本検証は限定的な模擬環境での実施であるため、ゼロデイ攻撃や迅速な初動対応の難しさを含む課題については、より多様な実環境での実証が必要である。正常行動の基準設定とその継続的なメンテナンスも、変化する環境に適応するうえで重要な要素であるが、その効果は環境依存の側面を持つ。

以上を踏まえ、ハイブリッド型のアプローチを実践する上では、運用の属人性低減策及び自動化技術の導入を推進することが重要である。また、正常行動のモニタリング基準を継続的に見直す仕組みを確立することで、検出精度の向上と運用効率の両立を目指す必要がある。

6 章 まとめ

6.1 組織における実践に向けた考察

脅威ハンティングは、その理論的有効性が広く認識されている一方で、実際の組織環境で効果的に実践し、継続的な活動として定着させるには複数の課題を乗り越える必要があると考えられる。本節では、第5章で詳述した実践検証から確認できた有効性や課題、そしてこれらの課題への対応策について考察する。

6.1.1 検証から見た有効性と実践上の課題

① 脅威ハンティングの有効性

脅威ハンティングは、既存のセキュリティ対策を補完し、組織のプロアクティブな防御能力を向上させる上で、以下のような有効性を示すと考えられる。

- 未知の脅威や正規ツールを悪用した攻撃に対する検出能力の向上

本検証では、正規の Windows 管理ツールなどを悪用する「LotL (Living-off-the-Land)」型の攻撃を再現し、攻撃ツールを一部カスタマイズすることで、一般的なアンチウイルスソフトによる検知を回避した。このような従来の防御をすり抜ける攻撃に対し、仮説を立てて異常なプロセスの挙動や不自然な通信の兆候に注目するアプローチをとりつつ、脅威ハンティングを実施した。その結果、定義ファイルやシグネチャベースの検知に頼らずに、侵害の痕跡を把握することができた。これは、従来型の検知技術では見落とされがちな攻撃に対しても、「検知の盲点」を補う手段として脅威ハンティングが有効であることを実証的に示すものである。

脅威ハンティングは、既存の対策や監視体制を代替するものではなく、それらを前提としたうえで行う補完的な取り組みであり、従来のセキュリティ体制に組み込むことで、防御の精度と深度を高める一助となると考えられる。

- セキュリティ運用上の課題の可視化と改善の契機

脅威ハンティングは、単に攻撃の兆候を見つけ出すための手段ではなく、既存のセキュリティ運用の課題を明らかにする契機にもなる。本検証では、仮説に基づいてログを確認・分析していく過程で、実際の運用上の問題点が浮かび上がった。

たとえば、調査に必要なログが一部のシステムで収集されていなかったり、取得できていても形式や内容が不統一であるケースがあり、こうした不備が脅威の検出精度に影響していた。また、既存の監視ルールでは想定されていない攻撃パターンも多く、脅威ハンティングを通じてそのギャップに気づくことができた。

このように、脅威を探索する活動そのものが、ログ収集の網羅性やルール設計の妥当性といったセキュリティ運用の状態を見直すきっかけとなる。

② 実践上の課題

有効性がある一方で、脅威ハンティングを組織で効果的に実践するには、以下のような課題に直面することが予想される。

- 高度な専門知識と攻撃手法への深い理解の必要性

脅威ハンティングは攻撃者の思考を読み解き、微細な痕跡から攻撃の兆候を捉えるため、脅威ハ

ンターは IT システムに精通していることを前提とし、多様な攻撃テクニックに関する高度な専門知識と深い理解が求められる。このような専門知識を持つ人材の不足は、脅威ハンティング実践の大きな障壁になると考えられる。

- ログ分析基盤の確立

脅威ハンティングを効果的に実施するためには、エンドポイントからクラウドまで多岐にわたるソースからの膨大なログデータの横断的分析の実現が求められる。しかし、現実には各ツールやデータがサイロ化しており、分析の障壁となることが少なくない。こうした状況を踏まえると、多様なツール群を効果的に連携させる分析基盤の構築も求められるが、その実現には多くの困難が伴うと予想される。

- ログの欠落

仮説に基づく脅威ハンティングを行う上では、必要なログデータの有無や品質が極めて重要となる。しかし現実には、そもそも検証に必要なログが収集されていなかったり、ログ自体は存在していても分析に必要な情報が欠落しているといったケースが少なくない。どれほど高度な分析基盤を整備していても、存在しないデータから脅威を発見することはできないため、適切なログの収集と内容の網羅性の確保が課題となり得る。

- 検出ロジックや検索クエリの網羅性・正当性の担保

脅威ハンティングでは、立てた仮説に基づいて、具体的な検出ロジックや検索クエリを作成する必要がある。これらが対象とする攻撃の痕跡を適切に捉えられるかどうか、また、過検出や誤検知を抑えつつ検出精度を維持できるかといった、網羅性と正確性の担保は、ハンティングの成果に直結する重要な要素である。

第 5 章でも示したように、仮説構築を行うためには MITRE ATT&CK のようなフレームワークを参考にしつつ、自組織の環境特性に応じたカスタマイズや継続的なチューニングが求められる。しかし、日々の運用に追われる中で、絶えず進化する攻撃者の TTPs を追いながら、膨大な検出ロジックや検索クエリを体系的に開発・評価・維持するには、高度なスキルを持つ人材はもちろんのこと十分な時間の確保が必要となる。そのため、網羅性を担保するような取り組みを日々の運用で継続的に確保していくことは難しいと考えられる。

- 明確な目的設定と適切な仮説構築の困難さ

第 4 章で詳述したように、脅威ハンティングは明確な目的設定、それに基づく検証可能な仮説の構築から始まる。しかし、「何を探すべきか」「どこに焦点を当てるべきか」を適切に定めて範囲を限定することは容易ではない。目的が曖昧であったり、仮説が広すぎたり、あるいは検証困難であったりすると、ハンティング活動は方向性を見失い、時間を費やしても成果が得られない「終わりのない作業」に陥る可能性がある。

6.1.2 組織への導入・定着に向けた要点

脅威ハンティングを組織に効果的に導入し、持続可能な業務プロセスとして定着させるためには、単に技術面の整備にとどまらず、組織体制や運用プロセスの観点からも包括的な取り組みが必要である。

本項では、ラトビアの国家情報セキュリティ機関 CERT.LV によって公開されている「Threat

Hunting Playbook」²¹の内容を参考としながら、本レポートにおける知見や検証結果を踏まえて、組織内で脅威ハンティングを定着させるために留意すべき要点を整理した。

① データ集約基盤の設計と構築

脅威ハンティングを実効的に行うためには、まず信頼性が高く、柔軟かつセキュアなデータ集約基盤の構築が不可欠である。データ集約基盤は単なるログの集約装置ではなく、調査の出発点となる分析基盤として設計原則を明確にして構築する必要がある。この際、将来的なデータ量や対象範囲の拡大に対応できる「拡張性」を備えた構成とすることも重要である。

実際の運用現場では、リソースなどの制約から、データ収集が計画通りに進まず、全体の構成を段階的に拡張・適応させながら構築するアプローチが現実的である。たとえば、初期段階では優先度の高いサーバー群に対象を限定し、後から範囲を広げる手法が挙げられる。

ログが存在すればハンティング自体は可能であるが、より多くのデータを収集することで検出精度が高まる一方、必要となる計算資源やストレージも増大する。さらに、脅威活動は発覚前から長期間にわたり潜伏している場合があり、過去数年分のデータが有効な手がかりとなることも少なくない。しかし、実際の運用では多くの組織が6～12か月程度のログ保存ポリシーを採用しており、これは主に技術的制約というより、容量やコスト、運用負荷とのバランスを取るための方針によるものである。したがって、基盤設計時には「何を、どこまで、どの期間保持するか」というトレードオフを見極めることが欠かせない。

また、ストレージ容量やネットワーク帯域、ログ収集エージェントの設置に必要なCPU・メモリリソースなどに加え、分析・可視化ツールが要求するリソースも見積もりに含め、システムに過剰な負荷をかけない設計が求められる。

さらに、データの収集と利用に際しては、セキュアなデータ交換チャネルの設計が必要である。ログデータやエンドポイントからの情報は、TLS等によって暗号化され、改ざんや漏えいが発生しないようにすることが必須である。また、これらのデータフローが経路上で監視や制御されていることも、組織のサイバーセキュリティを保つうえで欠かせない。

このようなデータ集約基盤を整備することで、組織のサイバーセキュリティを高めるとともに、自組織のインフラを保護しながらハンティング活動を展開できる。

② 可視性と確保とベースラインの定義

脅威ハンティングの成果を左右するのが、調査対象環境における「可視性」の確保である。可視性とは、システムやネットワーク上で何が起きているかを正しく把握することであり、そのためには、どの資産がどこにあり、どのようなログが、どの精度で収集されているかを明確にする必要がある。しかし、こうした情報は往々にして最新の状態で文書化されておらず、現場部門ごとの認識にもばらつきが見られる。このため、ネットワーク構成図、OS種別、ログの収集範囲や保持期間などの基礎情報は、関係部門へのヒアリングやアンケートといった人的アプローチによって収集することも重要である。こうした手法は、運用状況やビジネス上の背景を理解するうえで有効であり、検出結果の偽陽性を排除するための重要な文脈を提供する。

²¹ 「CERT.LV - Threat Hunt Playbook」 <https://cert.lv/en/threat-hunt-playbook>

加えて、実際のインフラ状況と文書との齟齬を埋めるために、資産管理ツールや監視エージェントを活用し、自動的かつ客観的にネットワークとホストの状態を把握する技術的アプローチも合わせて導入すべきである。こうして得られた情報を元に、「正常な状態（ベースライン）」を定義することが、異常検知や後続の調査のための土台となる。

③ 業務継続性への配慮

脅威ハンティングを実施する際には、基盤側から対象インフラに対してログの取得やクエリの実行といったアクセスが発生するため、その影響範囲とリソース負荷の見積もりについても事前に検討しておく必要がある。

本番システムを対象とする場合、CPU やディスク I/O、ネットワーク帯域などに与える負荷を適切に制御しなければ、業務に支障をきたす恐れがある。こうしたリスクを最小限に抑えるためには、対象環境へのアクセス方法や頻度、データの取得単位を戦略的に計画し、必要最小限のデータを効率的に取得する手段を確立することが求められる。特に大規模環境においては、全ての生ログをデータ集約基盤に転送するのではなく、必要な情報のみを収集・集約し、対象システム側の負荷を抑えることが望ましい。

加えて、負荷管理の観点からは、検証環境での試験や、調査ツールの動作におけるリソース使用の閾値制御（CPU 使用率が一定を超えた場合の自動中断など）といった仕組みの導入も有効である。こうした配慮を欠くと、システムの安定性や組織内での信頼を損ない、ハンティング活動の継続が困難になるリスクがある。

④ 調査アプローチの最適化

脅威ハンティングにおいて、最初に直面する課題のひとつは「どこから調査を始めるべきか」という点である。限られたリソースの中で膨大なデータを扱う必要があるため、闇雲に脅威を探索するのではなく、脅威インテリジェンスに基づく仮説立案と優先順位に基づいた調査アプローチが求められる。以下は、攻撃のフェーズにおいてどこを調査の焦点とするかについての思考過程の例である。

たとえば、多くの攻撃者は検知を避けるため、永続化のためのプログラムをネットワーク内のごく限られたホストにのみ設置する傾向がある。そのため、企業全体を横断してログや設定情報の発生頻度を比較することで、通常とは異なる低頻度の異常を浮かび上がらせることが可能となる。こうした頻度分析（Grouping and Counts）を活用すれば、ノイズに埋もれた微細な異常の検出が容易になり、効率的な調査を可能にする。

一方で、Initial Access（初期侵入）を起点とする調査は、あらゆる潜在的な侵入口に対して脆弱性や利用される手法を特定・分析する必要がある。これは攻撃パターンの多様性ゆえに非常に広範な検討が必要となり、網羅的な調査は困難となる。

また、Lateral Movement（水平移動）を起点とする方法は、ネットワーク構成やシステムの利用形態、ユーザーの業務習慣といった組織固有の情報に深く依存する。そのため、同じ手法を別のネットワークにそのまま適用することは難しく、再現性が低い。組織ごとに異なる分析が必要となることから、汎用的な調査手法としては扱いづらい。

さらに、攻撃者の目的が、特定ホストからのデータ窃取であった場合、情報を引き出すまでには永続化や C2 通信といった複数のフェーズを通過している可能性が高い。つまり、脅威を検出する

には、攻撃の最終目的そのものだけでなく、その目的に至る過程に着目することも重要と言える。

こうした背景を踏まえると、永続化を仮説の出発点とし、「頻度」と「逸脱」を軸に痕跡を抽出していくアプローチは、シンプルで再現性が高く、かつ実務的な負荷を抑えた調査を可能にする。構造的な仮説設定とデータ分析の組み合わせにより、脅威ハンティングの初動段階において、効率性と実効性のバランスを最適化することができる。

⑤ ツールとテンプレートによる調査の効率化と標準化

脅威ハンティングを継続的かつ効果的に実施するためには、調査プロセスの効率化と標準化が不可欠である。特に、仮説の立案から検証、棄却に至る一連の分析作業は属人化しやすく、判断のばらつきや作業の非効率を招く要因となりうる。

こうした課題に対処する手段として、過去に有効だった検索クエリや可視化手法、異常な挙動の分類基準などをテンプレートとして整理・蓄積することが挙げられる。これにより、新たな調査への着手が迅速化されると同時に、作業の再現性と品質の一貫性を確保できる。また、テンプレートは「思考の枠組み」としての役割も果たし、視野の偏りや不要な深掘りを防ぐ助けにもなる。

さらに、分析ツールの適切な選定と活用も、調査の精度と効率を大きく左右する。ログ検索、アノマリ検出、タイムライン構築など、目的に応じたツールを組み合わせ、操作性やツール間の連携性に配慮することで、アナリストの負荷を軽減しつつ、論理的な仮説検証を支援する環境が整う。

また、ハンティング活動を戦略的かつ効率的に進めるためには、活動全体を一元的に管理できる仕組みが有効である。たとえば Redmine のようなタスク管理ツールを活用し、調査案件をタスクとして登録・追跡する方法が挙げられる。各タスクには、仮説や分析結果、証拠を随時追加し、進捗に応じて「進行中」「保留中」「終了」「悪意あり」「疑わしい」「調査済み」といったタグや、「高・中・低」の優先度ラベルで分類する。こうした仕組みにより、アナリストは調査内容を体系的に記録・共有でき、ケースを発展させながら分析を深められる。チームリーダーはタスクの進行状況を俯瞰し、優先度の見直しやリソース配分、補足質問の提示を迅速に行える。結果として、作業分担が明確化され、重複作業の防止やレポート作成の効率化にもつながる。

さらに、限られたリソース環境下では、機械学習の活用が効果的である。たとえば、SIEM の機械学習・クラスタリング機能を用いれば、異常検出の初期段階を自動化し、迅速な対応を支援できる。これにより、システムログやネットワークデータ、メモリ・ディスクのアーティファクトなど、複数データソースからの異常検出が容易になる。もっとも、機械学習アルゴリズムはあくまでデータセットに基づいた「初期のヒント」を提示するに過ぎず、その結果は人間が確認すべきである。

最後に、生成 AI の活用についても言及しておきたい。現段階では、生成 AI を仮説立案の主軸として利用することは推奨されない。理由は二つある。第一に、生成 AI の出力は統計的パターン生成に基づいており、文脈理解や因果関係の把握には限界がある。そのため、専門知識や洞察を持つ人間が内容を検証・補完することが不可欠である。第二に、内部機密情報や組織固有のデータを外部環境に送信するリスクが存在することである。こうしたリスクを回避するためには、自社のシステムや管理された環境で生成 AI モデルを実行することが望ましい。その場合、生成 AI は仮説立案のための新たな視点や広い視野を提供し、補完的ツールとして活用できる可能性がある。

このように、テンプレートとツールを体系的に整備・活用することで限られた時間の中でどこにリソースを集中させるべきか、どの時点で仮説を棄却して方向転換すべきか、といった判断がしや

すくなり、脅威ハンティングのプロセスをより戦略的にコントロールできるようになる。結果として、調査精度の向上、作業の効率化、判断の明確化といった多面的な効果が得られる。

⑥ 成果の文書化とナレッジ共有による組織的活用

脅威ハンティングは、単に潜在する脅威を発見することだけが目的ではない。その過程や結果で得られた知見を文書化し、組織全体のセキュリティ体制の改善につなげることが重要である。最終的な報告書には、検出された脅威の有無に加え、活動過程で明らかになった設定ミス、監視の死角、運用ポリシーの不備など、組織内部に潜む構造的な課題も記載する。これらに対しては、優先順位づけされた改善提案を添えて文書化し、経営層から現場担当者に至るまで活用できるアウトプットとすることが望ましい。

また、ハンティング活動はプロジェクトとして一定期間にわたり実施されることもあるため、プロセス中の進捗や途中経過の文書化も重要である。中間報告として、仮説設定・検証状況、使用したデータソース、観測された異常の傾向などを適切に記録し、関係者との情報共有を図ることで、調査方向の見直しや支援リソースの調整にも活用できる。

さらに、活動で得られた知見は、担当者個人のノウハウにとどめず、組織の共有資産として体系的に蓄積・活用すべきである。そのためには、再利用可能なプレイブックの整備・更新が不可欠であり、検索クエリや可視化テンプレート、調査手順を整理し、事後のハンティング活動へと継承していく運用体制が求められる。

加えて、検出された新たな攻撃手法や行動パターン、脅威アクターの TTPs などの知見については、可能な範囲で匿名化・サニタイズを施しつつ、ISAC や業界横断的な脅威インテリジェンス共有コミュニティに対してフィードバックすることも有効である。これにより、自組織の知見が他者の防御に貢献するだけでなく、他組織から追加情報が得られる可能性もあるため、非常に重要な活動である。この際、データが第三者にとって実用的かつ有用であり続けるためには、脅威の性質や影響範囲、再現性の高い検出ロジックといった要素に適切なラベル付けを行い、文脈が損なわれないよう工夫することが求められる。

このように、脅威ハンティングの成果は、発見された事象の記録にとどまらず、調査過程や分析結果を通じて得られる「学び」を組織的な知識体系へと昇華させることに真価がある。これにより、単発の調査から継続的なセキュリティ強化サイクルへの移行が実現できる。

⑦ 持続可能な体制と評価文化の確立

脅威ハンティングを組織的な取り組みとして成功に導くためには、これまでに述べたようなデータ基盤や調査プロセス、ツールの整備に加え、それらを継続的かつ効果的に運用できる「組織体制」の構築が不可欠である。とりわけ、専門のハンティングチームが SOC や CSIRT、インフラ運用部門などと密接に連携し、部門横断的に情報と知見を共有できる体制が重要である。こうした連携があって初めて、より実効性の高いハンティング活動が可能となる。

また、活動の成果をどう評価するかという点も、脅威ハンティングの持続性を左右する重要な要素である。脅威ハンティングは必ずしも毎回、明確な脅威を発見できるとは限らない。したがって、「発見件数」や「アラート数」といった短期的・数量的な指標だけで成果を測定してしまうと、本来の意義が損なわれかねない。むしろ、潜在的なリスクの存在可能性を前提に、能動的に環境を探索し、攻撃の兆候を洗い出すという活動そのものの価値を、組織全体で正しく理解し評価す

る文化の醸成が求められる。

たとえ脅威が発見されなかったとしても、それは「現時点におけるシステムの健全性が確認できた」という重要な成果であり、継続的なセキュリティ評価プロセスの一環として意味のある結果である。こうした体制と文化を根付かせ、維持するためには、経営層の理解と支援が不可欠である。脅威ハンティングが単なる技術的試行ではなく、組織のリスクマネジメント戦略に直結する活動であるという認識を経営層が持ち、適切なリソースの継続的な配分と、活動の成果を中長期的視点で評価する姿勢を示すことが、持続可能な取り組みの鍵となる。

6.2 おわりに

本レポートでは、高度化・巧妙化するサイバー脅威に対抗するための積極的な取り組みとして、脅威ハンティングの重要性などについて述べてきた。本節では、これまでの内容を踏まえ、脅威ハンティングを实践するうえで核となる考え方や姿勢を改めて説明し、本レポートの締めくくりとする。

① 脅威を理解することの重要性

サイバー攻撃に対する有効な対策は、攻撃者が用いる手法や戦術を深く理解することなしには成り立たない。検証では、カスタマイズされたマルウェアが容易にセキュリティ製品を回避できることを確認し、こうした実態に即した仮説立案や検知設計の必要性を痛感した。このような高度かつ巧妙な攻撃に備えるためには、単に既存のルールやシグネチャに依存するのではなく、攻撃者の視点に立ち、そのテクニックを理解することによって、はじめて実効性のある具体的なセキュリティ対策を講じることが可能になると考えられる。

② 侵害前提思考に基づくセキュリティ対策の必要性

サイバー攻撃の高度化・複雑化が進む中で、「侵入されるかもしれない」ではなく「すでに侵入されているかもしれない」という「侵害前提」に立つことが、現実的なセキュリティ対策を実施するうえでの重要な考え方である。たとえば、マルウェアの侵入を防げなかったとしても、ネットワーク内で不自然な振る舞いをいち早く検出・対処することができれば、組織への影響を最小限に抑えることができる可能性がある。結果として「サイバーレジリエンス（回復力）」の強化にもつながり、組織の業務継続能力を向上させると考えられる。

③ 脅威ハンティングを支える組織づくりの重要性

脅威ハンティングは、単なる技術的な作業ではなく、組織としての継続的な取り組みとして設計・運用される必要がある。たとえ優れた人材がいたとしても、そのスキルや知識が属人化しては、ハンティング活動の再現性が担保されない。調査の手順や判断基準、分析結果の記録・共有といった運用基盤を整備することで、限られたリソースでもハンティング活動を円滑に進めることが可能となる。加えて、仮説立案・検証にかかる作業を効率化するテンプレートやツール群を整備することで、アナリストの思考をサポートし、調査の質を一定に保つことができる。

また、ハンティング活動は脅威が発見されなかったとしても、それは「現時点でのシステムの健全性が確認できた」という重要な成果であり、継続的なセキュリティ評価の一環として意義がある。こうした活動の成果を、「単に脅威があるか否か」ではなく、「環境全体のリスクを評価するプロセス」

として経営層に還元することで、組織としての理解と支援が得やすくなる。脅威ハンティングを持続的に実施していくためには、こうした評価軸の整理と、それに基づく文化の醸成が欠かせない。

④ アクティブサイバーディフェンスの必要性

脅威が高度化する今、セキュリティのあり方も「受動的な活動」から「積極的な活動」へと変化が求められている。日々届くアラートに反応するだけの運用では、未知の攻撃や巧妙な侵害に後手を踏む可能性が高まる。そのため、積極的にセキュリティ対策を講じる「アクティブサイバーディフェンス」の考え方が重要となる。

この考え方は、単に特定の技術や施策を指すものではなく、脅威を未然に探知・把握しようとする組織全体の姿勢そのものである。脅威ハンティングは、この姿勢を具現化する取り組みのひとつである。アクティブサイバーディフェンスは、従来のセキュリティ戦略の補完的な概念であり、組織文化の一部として定着させていくことが求められる。

謝辞

本レポート作成にあたり、産業サイバーセキュリティセンター中核人材育成プログラムの講師である、門林雄基先生、佐々木弘志先生には本レポートの元となる「脅威インテリジェンスを活用したセキュリティ強化のためのアプローチ」プロジェクトのメンター・講師としてご指導、ご助言とともにご支援を賜り続けてきました。改めて御礼申し上げます。

また、ラトビア共和国サイバーインシデント対応機関（CERT.LV）Dr.Bernhards・”BB”・Blumbersgs、その他、複数の脅威インテリジェンス関連企業の皆様にも、ヒアリングや本レポートのレビューなどの支援をいただきました。御礼申し上げます。

そして本レポートの作成や、本プロジェクトを共に実施した下記メンバーの皆様にも感謝を伝えたいと思います。

〈脅威インテリジェンスを活用したセキュリティ強化のためのアプローチ プロジェクトメンバー〉

【リーダー】

尾上将征

【サブリーダー】

大森育成

【メンバー】

川添恭平 木村光博 高橋勝 田川卓哉 田中貴道 綱川純平 増川京佑 山本大貴 和田歩

用語集

用語	説明
APT (Advanced Persistent Threat)	標的に対して長期間にわたり継続的に行われる高度な攻撃。国家や高度な犯罪集団が担うことが多い。
C2 (Command and Control)	攻撃者がマルウェアなどを介して侵害した端末を遠隔操作・制御するために使用する通信経路や仕組みのこと。
CISO (Chief Information Security Officer)	組織における情報セキュリティ全般の責任者。セキュリティ戦略や体制を統括する役職。
CSIRT (Computer Security Incident Response Team)	企業や団体内の情報セキュリティインシデントに対応する専門チーム。
ISAC (Information Sharing and Analysis Center)	同業他社で脅威情報を共有し、セキュリティ対応を連携するための業界横断的な組織。
IoA (Indicator of Attack)	攻撃そのものの兆候や行動に注目した検出指標。従来の IoC より早期対応が可能。
IoC (Indicator of Compromise)	攻撃による侵害の兆候を示す識別子。IP、ドメイン、ファイルハッシュなどが該当。
LotL (Living-off-the-Land)	正規ツール（例：PowerShell）を悪用して行われる、検知回避性の高い攻撃手法。
MITRE ATT&CK フレームワーク	MITRE が公開する、攻撃者の行動を戦術・技術として体系的に整理した知識ベース。
PESTLE 分析	政治 (Politics)、経済 (Economy)、社会 (Society)、技術 (Technology)、法制度 (Legal)、環境 (Environment) の 6 視点から脅威環境を分析する枠組み。
SOAR (Security Orchestration, Automation and Response)	セキュリティ運用の自動化・標準化を支援する基盤。インシデント対応の効率化に活用される。
SOC (Security Operation Center)	組織内外のセキュリティ監視・ログ分析・インシデント対応を 24 時間体制で行う拠点。
STIX/TAXII (Structured Threat Information Expression / Trusted Automated Exchange of Indicator Information)	脅威インテリジェンスを標準形式で表現・交換するための規格 (STIX) と、通信プロトコル (TAXII)。
Sysmon	Windows の詳細ログを収集するツールで、攻撃検出に活用される。
TTPs (Tactics, Techniques, and Procedures)	攻撃者が用いる戦術・技術・手法を指す。脅威分析の基本単位。
サイバー欺瞞	防御側が攻撃者を意図的に欺き、誤った情報や環境に誘導することで、攻撃の遅延・攪乱・発見を狙う防御戦術。ハニーポットや偽の認証情報、擬似システムなどを用いて、攻撃者の行動を観察・分析し、防御側の有利な情報を得ることを

	目的とする。
シグネチャベース	既知のマルウェアや攻撃に特有のパターン（シグネチャ）と一致するものを検出するセキュリティ手法。
ダークウェブ	通常の検索エンジンでは到達できない匿名性の高いインターネット領域。
ペイロード	攻撃で実行される不正なコードや処理内容。マルウェアや Exploit の一部。
脅威アクター	攻撃を実行する主体。国家、犯罪組織、ハクティビスト、内部不正者などが含まれる。