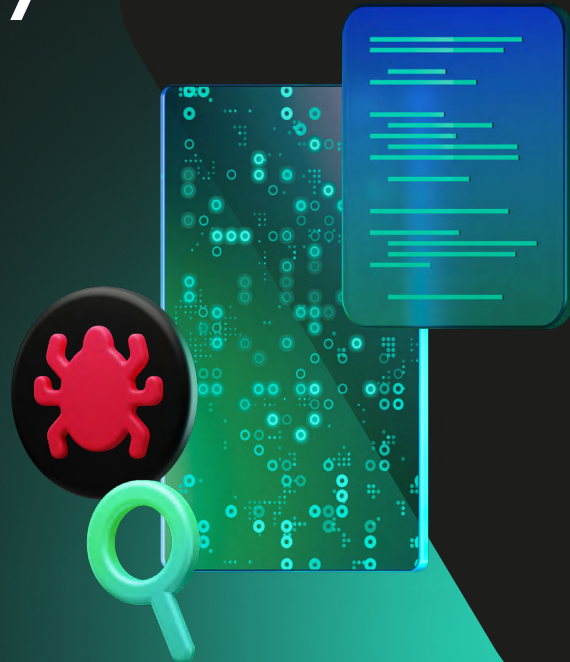




Поиск индикаторов компрометации на практике: эффективные техники выявления связанных индикаторов

Автор статьи – [Дамир Шайхелисламов](#), руководитель группы развития решений по защите от сложных угроз. В материале рассматривается роль анализа взаимосвязей между индикаторами компрометации в проактивном поиске угроз.

Автор демонстрирует, как на основе единственного исходного индикатора можно восстановить цепочку вредоносной активности и существенно расширить контекст расследования. Статья опирается на практические примеры, которые помогут оптимизировать процесс анализа угроз и пополнить базу индикаторов актуальными данными. Например:



исследование связанных индикаторов (IP-адресов, доменов, SSL/TLS-сертификатов)



обнаружение связанных артефактов вредоносного ПО на основе динамического и статического анализа



атрибуция угроз через анализ тактик, техник и процедур злоумышленников и их сопоставление с матрицей MITRE ATT&CK

На основе этих и других техник формируется системный подход, который помогает превращать отдельные индикаторы в полезную аналитическую информацию. Это позволяет аналитикам находить больше угроз, быстрее реагировать на инциденты и проактивно защищать инфраструктуру.

Поиск индикаторов компрометации и его роль

Современный ландшафт угроз не оставляет права на ошибку: иногда достаточно пропустить один индикатор, чтобы потерять из виду весь инцидент. Несмотря на развитие технологий поведенческого анализа, поиск индикаторов компрометации (IoC) по-прежнему остается одним из ключевых методов киберзащиты.

Индикаторы компрометации – это цифровые следы вредоносной активности: подозрительные IP-адреса, домены, URL, хеши файлов или изменения в реестре, которые могут служить отправной точкой для выявления более масштабной атаки.

Найти индикатор компрометации – это лишь первый шаг. Ключевая задача аналитика – выявить связанные с ним артефакты и понять, как они складываются в единую цепочку атаки. Такой анализ позволяет обнаружить скрытую инфраструктуру злоумышленников, проследить распространение угрозы и восстановить ход атаки на основе событий, которые на первый взгляд могут выглядеть разрозненными.

Выбор метода зависит от типа индикатора: их разнообразие велико, и охватить все возможные варианты в рамках одной статьи не представляется возможным. Поэтому ниже рассмотрим несколько практических подходов, которые аналитики применяют в повседневной работе.

Прежде всего – аналитика угроз

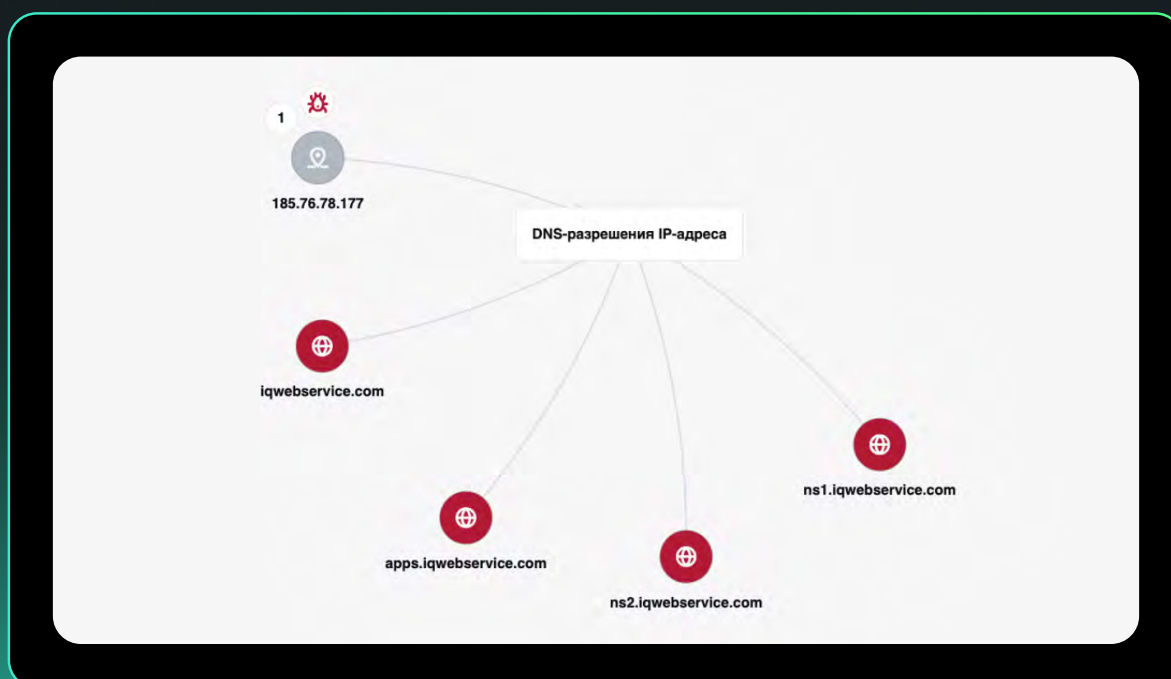
Исследование связанных индикаторов невозможно без качественной аналитики угроз: эффективность такого анализа напрямую зависит от полноты, актуальности и достоверности контекстных данных. Современные платформы аналитики угроз предоставляют инструменты для автоматизированного сопоставления, визуализации и обогащения данных, что существенно упрощает этот процесс.

Популярные техники анализа связанных индикаторов

Поиск доменов по IP-адресу

Чаще всего аналитикам приходится проверять подозрительные IP-адреса из оповещений системы безопасности. Предположим, система отправляет DNS-запросы на адрес **185.76.78.177**. Анализ данных passive DNS (pDNS) позволяет обнаружить домены, которые в прошлом были связаны с этим адресом.

Для оценки их релевантности следует проверить репутацию этих доменов через платформы аналитики угроз, а также поискать связанный с ними трафик в журналах DNS-запросов и прокси-сервера.



Граф pDNS-разрешений IP-адреса 185.76.78.177 – Kaspersky Threat Intelligence Portal

Поиск образцов вредоносного ПО по IP-адресу

Подозрительный IP-адрес, фигурирующий в оповещениях и журналах сетевой активности, нередко требует дальнейшей проверки. Анализ связанных индикаторов позволяет определить, в какой вредоносной активности задействован данный IP-адрес:

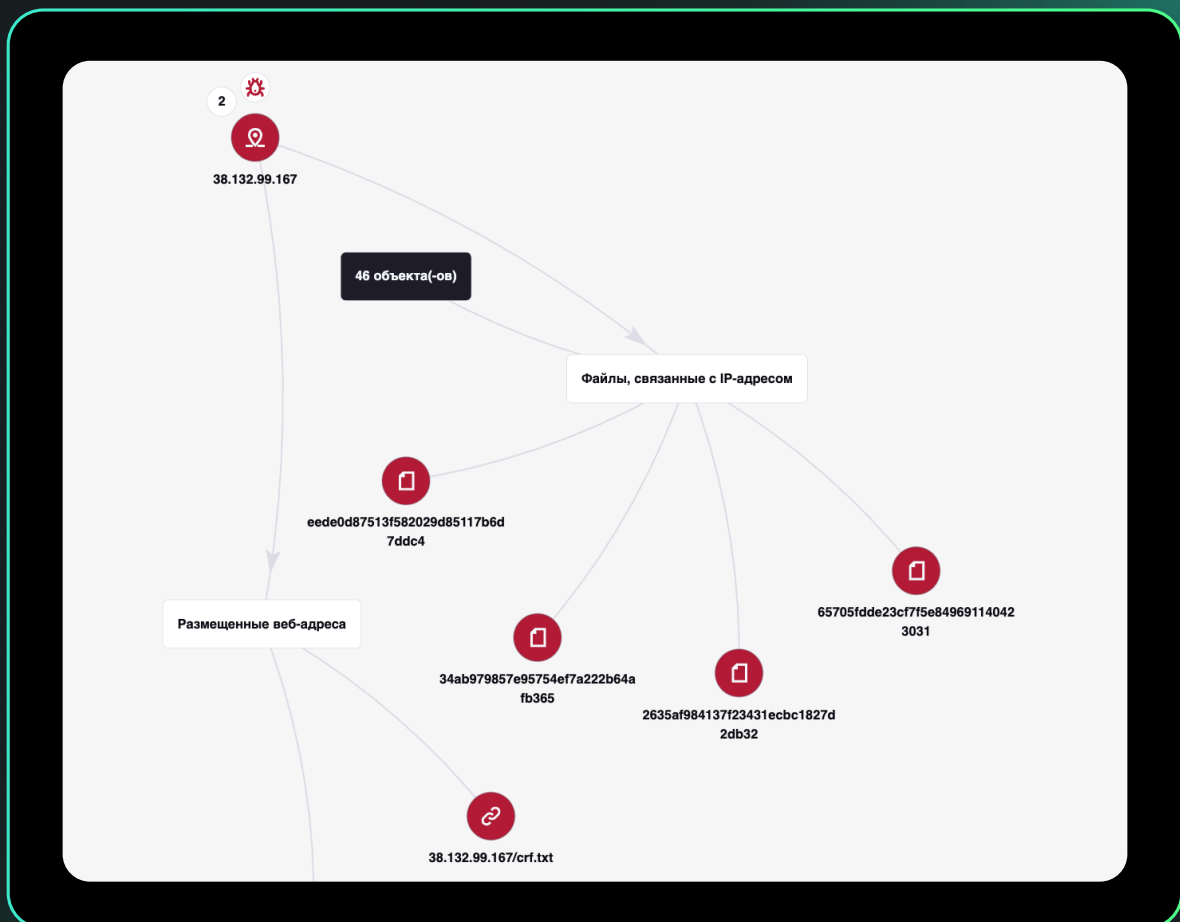
1

Размещение образцов вредоносного ПО.

2

Использование в качестве сервера управления и контроля (C2).

Как общедоступные, так и коммерческие платформы аналитики угроз предоставляют широкие возможности для сопоставления IP-адресов с известной вредоносной активностью и связанными кампаниями.



Хеши файлов, связанных с IP-адресом 38.132.99.167 – [Kaspersky Threat Intelligence Portal](#)

SSL/TLS-сертификаты

Сертификаты SSL/TLS служат важными опорными точками для активного поиска угроз и анализа связанной инфраструктуры. Злоумышленники нередко используют одни и те же сертификаты для разных серверов и доменов, входящих в их инфраструктуру, а также активно применяют бесплатные центры сертификации, например Let's Encrypt, автоматически выпускающие сертификаты с минимальной проверкой.

Злоумышленники нередко повторно используют схожие шаблоны и параметры при выпуске сертификатов. Ниже приведен пример типичной структуры, которая может встречаться в рамках нескольких вредоносных кампаний: **C=US, ST=California, L=San Francisco, O=Microsoft Corporation, OU=Security Division, CN=(доменное имя).**

Это открывает широкие возможности для анализа связанной инфраструктуры: аналитики могут выявлять IP-адреса ранее использовавшихся серверов, на которых были установлены сертификаты с такими же или схожими данными о субъекте.

Помимо сертификатов, для анализа связанных индикаторов и кластеризации инфраструктуры злоумышленников могут применяться метаданные TLS-рукопожатий, включая хеши JA3, JA3S и JARM. Они представляют собой уникальные цифровые отпечатки, отражающие особенности обмена данными между клиентом и сервером по протоколу TLS. Злоумышленники нередко используют одинаковые конфигурации или вредоносные фреймворки на разных серверах, что порождает идентичные или почти идентичные отпечатки.

Например, JA3-сигнатура **b742b407517bac9536a77a7b0fee28e9** ассоциируется с фреймворком командного сервера Cobalt Strike. Сопоставив этот отпечаток с данными сетевой телеметрии или наборами аналитических данных, можно найти и другие вредоносные серверы этой группировки.

Домен → TXT-запись DNS

При поиске индикаторов компрометации расследование часто начинается с проверки доменов. Однако анализ связанных с ними TXT-записей DNS – не менее эффективный способ выявления дополнительного контекста, который нередко остается без внимания.

TXT-записи могут использоваться злоумышленниками для хранения служебных данных, резервных адресов C2-инфраструктуры, токенов, ключей или фрагментов полезной нагрузки. Проследив связь от подозрительного домена к соответствующим TXT-записям, аналитики могут:



выявить новые домены и резервные адреса командных серверов



обнаружить команды или конфигурационные данные ботнета



извлечь ключи шифрования или токены



выявить признаки злоупотребления протоколом DNS, включая туннелирование или скрытую доставку полезной нагрузки



восстановить отдельные фрагменты файлов или компонентов полезной нагрузки

Например, при проверке фишингового домена была обнаружена TXT-запись, содержащая IP-адреса резервных командных серверов, закодированные в формате Base64. Располагая этими данными, SOC может заблаговременно заблокировать их.

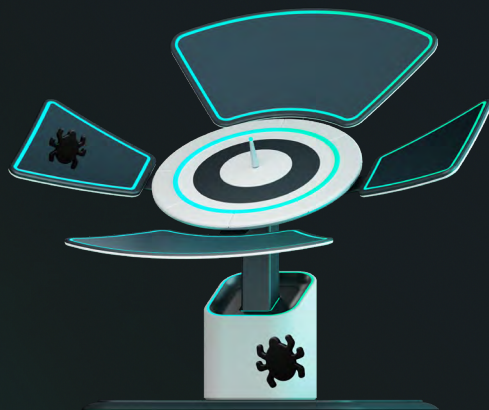
Временные интервалы

Периодические обращения скомпрометированных узлов к командному серверу выполняют роль сетевых маячков (beacons). Такие сигналы подтверждают, что зараженная система остается активной и доступной для управления, а злоумышленник сохраняет контроль над атакой.

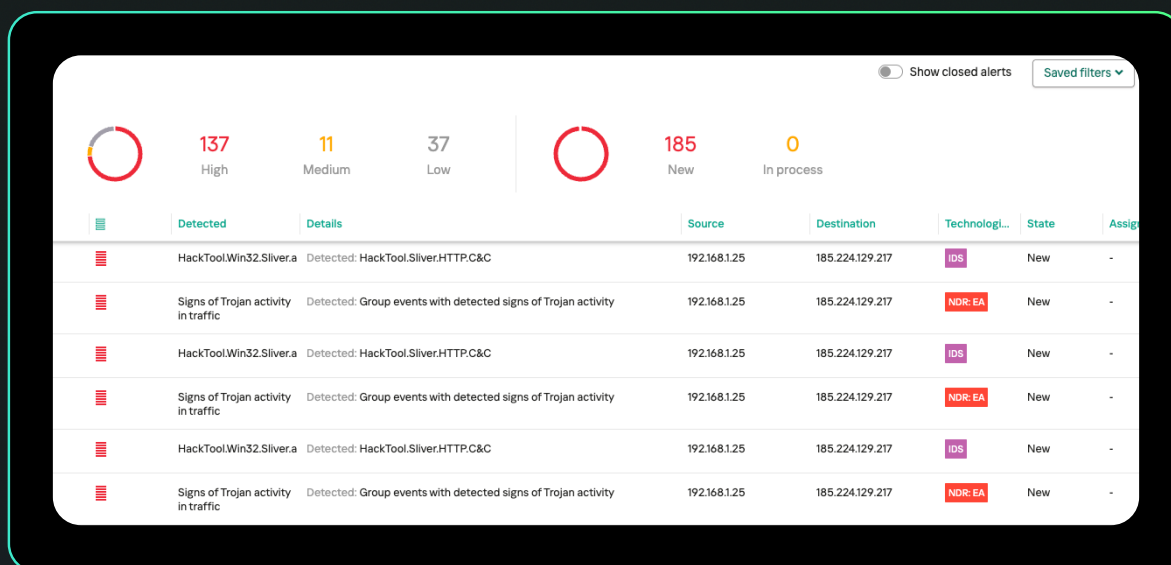
Как правило, такие обращения происходят через равные или почти равные промежутки времени – например, каждые несколько минут. Именно эта регулярность отличает их от естественной вариативности сетевого трафика.

Чтобы затруднить обнаружение, злоумышленники могут намеренно вносить случайные отклонения в интервалы между обращениями. Тем не менее выявить подозрительные закономерности позволяет статистический анализ: в частности, оценка среднего интервала между обращениями и стандартного отклонения.

Анализ регулярности, средних интервалов и случайных отклонений помогает обнаруживать аномалии, сопоставлять их с журналами событий и выявлять новые зараженные узлы или связанные образцы вредоносного ПО.



Например, в ходе расследования инцидента аналитики обнаружили исходящие сетевые подключения, устанавливаемые через равные промежутки времени, – поведение, характерное для маячка командного сервера Sliver.



Регулярные подключения к командному серверу Sliver – **Kaspersky Anti Targeted Attack (NTA)**

От анализа кода к расширению базы индикаторов компрометации

Каждый бинарный файл содержит характерные особенности кода – своего рода «генетические маркеры», которые могут указывать на связь с известной целевой атакой или АРТ-кампанией. Методы атрибуции угроз, в том числе с использованием решения [Kaspersky Threat Attribution Engine](#), позволяют аналитикам сопоставлять фрагменты обнаруженного кода с репозиториями образцов сложного вредоносного ПО.

Чтобы извлечь максимум пользы из этой информации, аналитики могут использовать фрагменты кода для создания правил для YARA, которые в дальнейшем позволят автоматически обнаруживать схожие бинарные файлы в корпоративных системах, репозиториях вредоносного кода или песочнице. В процессе анализа могут быть выявлены новые индикаторы компрометации, способствующие обнаружению развернутых в инфраструктуре вредоносных программ или следов более ранних атак.

В данном случае применяется метод атрибуции кибератак на основе анализа кода: анализ уникальных характеристик бинарных файлов позволяет соотнести их с конкретной АРТ-кампанией, упрощая выявление связанных индикаторов компрометации и реконструкцию хронологии прошлых атак.

Анализ артефактов, полученных в песочнице

Часть индикаторов компрометации выявляется только при анализе подозрительного объекта в контролируемой среде. Многие злоумышленники проводят бесфайловые или многоступенчатые атаки: начальная полезная нагрузка может быть безвредной или минимальной, тогда как полноценная вредоносная активность разворачивается лишь после запуска доставленных в систему объектов.

Выполняя подозрительный код в песочнице, аналитики могут исследовать его поведение в среде выполнения – исходящий трафик, внедряемые файлы, изменения в реестре и активность процессов, – чего не позволяет достичь статический анализ. На основе этих поведенческих артефактов формируются новые индикаторы компрометации, обогащающие аналитические базы и повышающие качество обнаружения угроз.

Хеш файла или URL-адрес могут служить отправной точкой для обнаружения других артефактов. Поиск по хешу (MD5, SHA1, SHA256) или URL-адресу – одна из наиболее эффективных техник, широко применяемых в анализе угроз.

Запустив вредоносный файл в песочнице, можно установить, какие объекты он внедряет в систему: исполняемые файлы, DLL-библиотеки, скрипты, файлы конфигурации – все они являются компонентами многоступенчатых атак.



Результаты

Действия в системе

Извлеченные файлы

Действие в сети

Файлы, извлеченные из объекта

Скачать данные

Статус	Категория	MD5	Имя файла	Уникающие	Тип файла	Детектируемые объекты
Вредоносный	—	0a7b32d23fbd6d62a593c234bafa2311	e9535d0d5e8e17779b496d7988c0b0547efbdaabb482da8497a5f0da87cbefc96	ZIP	Microsoft Word doc	HEUR: Trojan.Win32.Generic; Trojan.MSWord.Agent.m

Файлы, извлеченные из сетевого трафика

Нет данных

Сохраненные файлы

Скачать данные

Статус	Категория	MD5	Детектируемый объект	Имя файла	Тип
Вредоносный	ДРТ	0e7b32d23fbd6d62a593c234bafa2311	HEUR: Trojan.Win32.Generic	e9535d0d5e8e17779b496d7988c0b0547efbdaabb482da8497a5f0da87cbefc96.doc	Microsoft Word doc
Не определено	—	05f4f020520294c7a23793b308a7c8	—	index.dat	text
Не определено	—	1280460958a2c329f940223a8fca09	—	index.dat	text
Не определено	—	636c38a6b6c9c9f8a792a70834565	—	-\$356d0d5e8e17779b496d7988c0b0547efbdaabb482da8497a5f0da87cbefc96.doc	unknown
Не определено	—	94339c340c283977f3c03c0d88b403c	—	MSForms.exe	binary/tb
Не определено	—	ca4a81e34e8a5408a4528e9e0d0ca	—	EF079929.wmf	unknown
Не определено	—	c8a607073f92447c0e34f7b0a480a	—	407f9b08.wmf	unknown
Не определено	—	c5b87c137c3e4a7f4e3423f3c9b470	—	e9535d0d5e8e17779b496d7988c0b0547efbdaabb482da8497a5f0da87cbefc96.doc.LNK	lnk
Не определено	—	d046d0b720f70a4c96b8c0204f0f3	—	~W99(0284533-4229-4C09-8B18-CE7C80D28E8).tmp	unknown
Безопасный	—	A23C9387786030A084C8BAC2836	—	WMIC.exe	exe x32

Хеши файлов, внедренных анализируемым образцом (**0e7b32d23fbd6d62a593c234bafa2311**) – [Kaspersky Threat Intelligence Portal](#)

Зная, какие объекты внедрены в систему, аналитики могут оценить масштаб инцидента, реконструировать цепочку атаки и создать правила для YARA на основе сигнатур, имен и шаблонов поведения внедренных файлов.

Сопоставление ключей реестра с тактиками и методами злоумышленников (MITRE ATT&CK)

Изменения в реестре, которые ослабляют защиту системы или конфиденциальных данных, могут служить информативным индикатором компрометации. Характерный пример – вредоносное изменение параметров аутентификации WDigest, вследствие которого учетные данные сохраняются в системной памяти в открытом виде.

Столь незначительная, на первый взгляд, модификация способна существенно расширить поверхность атаки, предоставив злоумышленникам возможность извлекать учетные данные непосредственно из памяти процесса службы проверки подлинности локальной системы безопасности (LSASS).

Для снижения этого риска в современных версиях Windows протокол **WDigest** по умолчанию отключен (**UseLogonCredential = 0**). Но если злоумышленник изменит значение ключа на 1, система начнет кешировать учетные данные пользователей в памяти, откуда они могут быть извлечены специализированными инструментами для получения учетных данных из памяти.

Индикаторы компрометации, связанные с изменением реестра, помогают установить цели и техники злоумышленника.

Для сопоставления этих индикаторов компрометации с известным поведением злоумышленников можно использовать MITRE ATT&CK Navigator. Так, индикатор, свидетельствующий о попытке доступа к учетным данным – значение **UseLogonCredential = 1** в ветке WDigest, соответствует технике T1003.001 – OS Credential Dumping: LSASS Memory.

Такие инструменты, как **EDR** и **SIEM**, могут эффективно использоваться для обнаружения индикаторов компрометации и их сопоставления с тактиками злоумышленников.



Эффективный поиск индикаторов с помощью ИИ

Эффективность анализа связанных индикаторов всегда зависела от опыта аналитиков, однако современные ИИ-технологии позволяют существенно масштабировать этот процесс.



Автоматическое сопоставление. Модели машинного обучения за секунды находят связанные индикаторы компрометации в огромных наборах данных, даже при наличии незначительных различий между ними, таких как вариации поддоменов или аномалии в сертификатах.



Выявление закономерностей. ИИ способен обнаруживать повторяющиеся паттерны инфраструктуры, которые могут остаться незамеченными при ручном поиске угроз.



Фильтрация нерелевантных данных. Интеллектуальные модели отсеивают ложноположительные срабатывания, позволяя аналитикам сосредоточиться на критичных оповещениях.

Сочетание экспертной валидации со стороны человека и возможностей ИИ обеспечивает более быстрое и точное обнаружение индикаторов компрометации.

Пример из практики

Анализ даже одного хеша может помочь выявить полную цепочку заражения, обеспечивая как сдерживание инцидента, так и обогащение набора индикаторов компрометации.

Предположим, команда центра мониторинга безопасности обнаруживает подозрительный исполняемый файл (agent.exe) на конечном устройстве в финансовом отделе. EDR-система помечает уникальный хеш этого файла как подозрительный. Аналитики загружают хеш на платформы аналитики угроз, которые связывают его с известным семейством вредоносного ПО, специализирующегося на краже данных и DNS-туннелировании.

Отчет платформы содержит связанные хеши, ассоциированные домены (например, corp-updates[.]com), а также сведения о предыдущих кампаниях, в которых DNS использовался для эксфильтрации данных.

Опираясь на обогащенные индикаторы, полученные в результате анализа хеша, аналитики проверяют журналы DNS-запросов на наличие подозрительных доменов и длинных рандомизированных поддоменов, связанных с corp-updates[.]com. В ходе анализа они обнаруживают сотни исходящих DNS-запросов с нескольких конечных устройств к длинным случайно сгенерированным поддоменам corp-updates[.]com.

Подобный паттерн характерен для известных техник DNS-эксфильтрации, при которых вредоносное ПО кодирует похищенные данные непосредственно в поддоменах DNS-запросов.

Аналитики реконструируют цепочку атаки:



В результате целевой фишинговой атаки на устройство попал вредоносный документ, который загрузил и запустил файл agent.exe.



После закрепления в системе агент собирал конфиденциальные данные и передавал их через DNS-запросы. Такой канал сложнее обнаружить, поскольку он может маскироваться под легитимный DNS-трафик.



Анализ связанных хешей и индикаторов компрометации позволил выявить другие зараженные компьютеры, которые обращались к поддоменам helper[.]corp-updates[.]com и auditlogs[.]corp-backups[.]com.

Центр мониторинга безопасности блокирует на уровне межсетевого экрана все исходящие запросы к corp-updates[.]com и связанным доменам, а все устройства с совпадающими хешами или подозрительными DNS-запросами изолируются.

Заключение

Работа с индикаторами компрометации – это циклический процесс: каждый новый ИОС может привести к другим артефактам и расширить контекст расследования. Для такого анализа используются как статические индикаторы, например хеши файлов, IP-адреса и домены, так и динамические артефакты, выявленные в песочнице. Эти данные помогают аналитикам проследить инфраструктуру злоумышленников, находить связанные заражения и обогащать внутреннюю телеметрию.

Практическое задание. Выберите индикатор компрометации из одного из недавних расследований – в идеале связанный с подтвержденной вредоносной активностью – и проведите анализ как минимум по трем направлениям: инфраструктура (пассивный DNS, хеши JA3/JA3S), артефакты вредоносного ПО (файлы, внедренные в ходе анализа в песочнице, срабатывания правил YARA) и поведение злоумышленников (тактики и методы, соответствующие матрице MITRE ATT&CK).

Несмотря на все преимущества, работа с индикаторами имеет и ряд ограничений:



Высокая скорость смены индикаторов.

Индикаторы, в частности, домены командных серверов, нередко обновляются с высокой частотой: инфраструктура злоумышленников способна исчезнуть за считанные часы.



Ложноположительные срабатывания.

Возникают при использовании общедоступных потоков данных без дополнительной верификации.



Избыток информации.

Анализ по индикаторам способен генерировать значительный объем данных, требующих последующей обработки и приоритизации.



Техники обхода защиты.

Опытные злоумышленники все активнее применяют шифрование при доставке полезной нагрузки, сокрытие трафика за легитимными доменами (domain fronting) и обфускацию кода в среде выполнения.

Таким образом, анализ по индикаторам компрометации существенно повышает эффективность обнаружения угроз и реагирования на них, однако этот метод должен применяться в связке с поведенческим анализом, обнаружением аномалий и проактивным поиском угроз на основе тактик и методов злоумышленников.

О «Лаборатории Касперского»

«Лаборатория Касперского» – международная компания, работающая в сфере информационной безопасности и защиты конфиденциальности данных с 1997 года. Глубокие экспертные знания и многолетний опыт компании лежат в основе защитных решений и сервисов нового поколения, обеспечивающих безопасность бизнеса, критически важных инфраструктур, государственных учреждений и рядовых пользователей, защищая миллиарды устройств от новейших угроз и целевых атак. Обширное портфолио «Лаборатории Касперского» включает в себя передовые продукты для защиты рабочих мест, а также ряд специализированных решений и сервисов для борьбы с комплексными и постоянно эволюционирующими киберугрозами, в том числе кибериммунные системы. Наши технологии обеспечивают защиту самых важных аспектов бизнеса более 200 тысяч корпоративных клиентов. Подробнее см. на сайте www.kaspersky.ru.