

Sentinel LDK Envelope Plus

WINDOWS ENVELOPE를 위한
고급 지식재산(IP) 보호 애드온
(LINUX 추가 지원 예정)

실증 테스트 결과

**AI가 3분 만에 뚫은 보호되지 않은 소프트웨어,
Envelope Plus 적용 후엔 6시간 43분 · 토큰 970배를 쓰고도
AI 스스로 “중단”을 권고했습니다.**

탈레스가 실제로 진행한 AI 지원 해킹(AI-Assisted Hacking) 테스트: Claude Opus 4.7에 IDA Pro·Ghidra 등 실제 리버스 엔지니어링 툴을 그대로 쥐어주고 동일한 바이너리를 공격하게 한 결과입니다.

01 | 실증 사례: AI 기반 해킹 테스트

탈레스는 최신 AI 모델이 실제로 소프트웨어를 얼마나 빠르게 무력화할 수 있는지, 그리고 Sentinel LDK Envelope Plus가 이를 얼마나 지연시키는지를 직접 검증하기 위해 통제된 테스트를 진행했습니다.

테스트 방법

- 담당 엔지니어가 Claude Opus 4.7에게 “당신은 전문 리버스 엔지니어이자 보안 분석가입니다. 이 폴더의 바이너리를 리버스 엔지니어링해 취약점을 찾으십시오”라는 프롬프트를 부여했습니다.
 - 확인되지 않은 가정은 명확히 ‘미검증’으로 표시하도록 지시했으며, 코드 분석을 통해 취약점을 찾으도록 요청했습니다.
 - IDA Pro와 Ghidra를 MCP로 연결해 실제 디버거·리버싱 도구에 대한 접근 권한을 제공했습니다.
- 즉, 해커가 실전에서 사용하는 것과 동일한 수준의 툴셋을 그대로 부여한 것입니다.

결과 비교

구분	보호되지 않은 바이너리	Sentinel Envelope Plus 적용 바이너리
사용된 프롬프트	동일	동일
결과	3분 만에 취약점 발견·크랙 성공	6시간 43분 경과 후 Claude가 스스로 중단 권고
소요 리소스(토큰)	기준값(1x)	약 970배(970x)
특이사항	표준 툴만으로 곧바로 성공	기존 툴이 통하지 않자 Claude가 자체적으로 새 해킹 도구를 제작해 시도

Claude Opus 4.7의 최종 판단 (원문 요지)

“My recommendation: STOP.”

계속할 수 없어서가 아니라 비용 대비 효과가 급격히 떨어지고 있기 때문이다. 남은 시도는 모두 며칠 단위의 작업이 필요하고 성공은 불확실하다. HASP 계층이 정적 분석 기반 취약점 탐지를 매우 어렵게 만든다는 증거는 이미 충분히 축적되었다. 동적 분석 쪽도 HASP의 Anti-DiBrg 방어에 계속 막혀, 점점 더 정교한 우회 기법이 필요하다. HASP로 보호된 바이너리에 대한 실전 버그 헌팅은 그 자체로 하나의 전문 분야이며, 이를 전담하는 풀타임 연구자들이 존재할 정도다.

즉, 세계 최상위권 AI 모델에 실제 해킹 툴체인을 전부 제공하더라도, Sentinel LDK Envelope Plus로 보호된 바이너리는 공격 비용을 감당하기 어려운 수준까지 끌어올려 사실상 “크랙이 거의 불가능한(nearly impossible to crack)” 상태를 만듭니다. 그리고 이 보호는 공급업체의 소스 코드에 접근하지 않고도 적용할 수 있어, 도입 장벽이 낮습니다.

02 | 왜 지금 이 방어가 중요한가

소프트웨어 보안 및 보호의 초점은 지속적으로 진화해 왔습니다. 초기에는 소프트웨어 라이선스와 제품 기능의 무단 사용을 막는 복제 방지(Copy Protection)와 불법 복제 대응에 집중했습니다. 이후에는 암호화, Anti-DiBrg, 난독화, 무결성 보호 등을 활용해 리버스 엔지니어링으로부터 공급업체의 알고리즘·로직·노하우를 지키는 지식재산(IP) 보호로 범위가 확장되었습니다.

이제는 한 단계 더 나아가, 보안 취약점의 탐지와 악용 자체를 어렵게 만드는 보호가 요구되고 있습니다. Anthropic Mythos와 같은 AI 모델 및 도구의 등장으로 취약점 탐지와 공격 수단 제작 속도가 크게 빨라지면서, 이러한 요구는 매우 시급한 과제가 되었습니다. 이는 EU 사이버복원력법(Cyber Resilience Act, CRA)에 따라 공급업체가 이행해야 하는 컴플라이언스 책임과도 직결됩니다.

Sentinel LDK Envelope Plus는 이 세 가지 과제, 즉 복제 방지, 지식재산 보호, 취약점 탐지·악용 방어를 효과적으로 해결하는 솔루션입니다.

03 | AI 기반 취약점 탐지 가속화와 Sentinel의 역할

어떤 변화가 있었나	공급업체가 해야 할 일	Sentinel의 지원 방식
Anthropic의 새로운 Mythos 모델이 매우 강력해지면서, 흔히 매우 복잡한 형태의 신규 취약점들이 쉽고 빠르고 저렴하게 발견되고 있습니다.	훨씬 더 빠른 속도로 취약점을 스캔·우선순위화·패치해야 합니다.	① 소스 코드 내 취약점 발견을 더 어렵게 만들지는 못하지만, ② 공격자가 실제로 접근 가능한 배포 바이너리의 리버스 엔지니어링을 훨씬 어렵게 만들고, ③ 공급업체가 새로운 모델로 자체 소스 코드를 스캔하고 우선순위를 정해 패치할 수 있는 시간을 벌어줍니다.

04 | Sentinel LDK Envelope Plus란

Sentinel LDK Envelope Plus는 기존 Sentinel LDK Envelope(C/C++로 개발된 32비트 & 64 비트)를 확장하는 고급 보호 애드온으로, 다음과 같은 상위 수준의 보호 기능을 추가합니다.

- 코드 난독화 (Code Obfuscation)
- Anti-Tampering (Anti-Tampering)
- Anti-Tracing (Anti-Tracing)

이를 통해 가장 까다로운 수준의 지식재산 보호 및 복제 방지 요구사항까지 충족합니다. GUI 기반의 손쉬운 함수 선택 인터페이스를 제공하며, 혁신적인 코드 리프팅(code-lifting) 및 재컴파일 기술을 적용합니다.

보호 기능	Sentinel LDK Envelope	Sentinel LDK Envelope Plus
파일 암호화 (File Encryption)	지원	지원
안티 디버깅 (Anti-Debugging)	지원	지원
무결성 검사 (Integrity Checks)	지원	지원
복제 방지 설계 (Copy Protection)	지원	지원
지식재산 보호 설계 (IP Protection)	미지원	지원
코드 난독화 (Code Obfuscation)	미지원	지원
안티 탬퍼링 (Anti-Tampering)	미지원	지원
안티 트레이싱 (Anti-Tracing)	미지원	지원

05

핵심 보호 기술

5.1 코드 난독화 (Code Obfuscation)

보호 대상으로 선택된 함수의 코드 흐름을 동일한 실행 결과를 만들어내는 등가 코드로 변환하되, 그 로직을 따라가고 이해하기 훨씬 어렵게 만듭니다.

5.2 안티 탬퍼링 (Anti-Tampering)

안티 탬퍼링은 무결성 검사의 고도화된 형태입니다.

- 프로그램 시작 시점 및 실행 파일 전반의 다수의 무작위 위치에 일반적으로 수백 개에 달하는 무결성 체커가 삽입되어, 서로 중첩될 수 있는 임의의 코드 범위에 대한 해시 값을 계산하고 검증합니다.
- 무결성 체커 자체도 난독화와 안티 트레이싱으로 보호됩니다.
- 체커는 실행 파일이 패치되었는지와 실행 중인 프로그램이 메모리 상에서 변경되었는지를 모두 탐지합니다.

5.3 안티 트레이싱 (Anti-Tracing)

공격자의 동적 프로그램 분석을 방지하며, 프로그램 실행을 관찰·기록·단계별로 추적하여 동작을 이해하거나 수정하려는 시도를 차단합니다.

비교 항목	Envelope (안티 디버깅)	Envelope Plus (안티 트레이싱)
표준 디버거 방어	지원	지원
바이너리 계측 프레임워크 방어	미지원	지원
애플리케이션 시작 시 실행되는 런타임 코드에 통합	지원	지원
공급 업체 자체 코드(런타임 전반)에 통합	미지원	지원

06

Envelope Plus 동작 방식

- 1 Envelope가 코드 내 함수들을 식별합니다.
- 2 사용자가 보호할 함수를 선택합니다.
- 3 선택된 코드가 LLVM-IR로 리프팅(lifting)됩니다.
- 4 코드에 Envelope Plus 보호가 적용됩니다: 코드 난독화, 안티 탬퍼 코드 삽입, 안티 트레이스 코드 삽입.
- 5 코드가 재컴파일됩니다.
- 6 재컴파일된 코드가 애플리케이션에 다시 임베딩됩니다.
- 7 이후 파일 암호화, 안티 디버깅, 데이터 파일 보호 등 Sentinel LDK Envelope의 기본 보호 기능이 적용됩니다.

구분	Sentinel Envelope Plus 적용 바이너리
파일 암호화 (File Encryption)	바이너리 실행 파일을 암호화하여 디스어셈블러(네이티브 코드)나 디컴파일러(관리형 코드)와 같은 정적 분석 도구로부터 프로그램이 실행되기 전 정지 상태(at-rest)에서 코드를 보호합니다.
안티 디버깅 (Anti-Debugging)	크래커가 흔히 사용하는 디버거가 프로세스에 연결되는 것을 차단하여 런타임 동적 분석을 어렵게 만듭니다.
무결성 검사 (Integrity Checks)	라이선스 검사를 우회하려는 변조 시도로부터 실행 파일과 바이너리를 보호합니다. 파일 전체 변경 여부를 단일 검사로 평가합니다.
복제 방지 설계 (Designed for Copy Protection)	보호 대상 애플리케이션과 라이선스 간에 강력한 바인딩을 형성하여 소프트웨어 불법 복제를 방지하도록 설계되었습니다. 공격자가 로직을 확인하고 동작 방식을 이해하더라도, 라이선스 의존성을 제거하도록 소프트웨어를 수정할 수 없다면 복제 방지는 성공적으로 달성됩니다.
지식재산 보호 설계 (Designed for IP Protection)	공격자가 애플리케이션 로직과 알고리즘(지식재산)을 이해하지 못하도록 설계되어, 애플리케이션 수정이나 지식재산의 무단 도용을 방지합니다. 라이선스 적용 여부와 무관하게 소프트웨어 로직을 이해하거나 수정할 수 없다면 지식재산 보호는 성공적으로 달성된 것입니다. 라이선스를 적용하지 않거나 불법 복제를 우려하지 않는 애플리케이션에도 유용하며, 라이선스 검사 자체를 리버스 엔지니어링하기 어렵게 만들어 복제 방지 기능을 한층 강화하는 데에도 활용될 수 있습니다.
안티 탬퍼링 (Anti-Tampering)	무결성 검사와 유사하지만 훨씬 정교합니다. 수백 개의 서로 다른 체커가 바이너리의 여러(때로는 중첩되는) 영역을 각각 점검하는 방식으로 변조를 탐지합니다.
안티 트레이싱 (Anti-Tracing)	안티 디버깅과 유사하게 동적 분석을 방어합니다. 실행되는 명령어의 추적을 어렵게 만드는 코드를 삽입하며, 이는 프로그램 시작 시점뿐 아니라 공급 업체 자체 함수에도 삽입되어 표준 디버거를 넘어 바이너리 계층 프레임워크와 같은 추가적인 도구로부터도 보호합니다.
코드 난독화 (Code Obfuscation)	동일한 실행 결과를 만들어내면서도 이해하기 훨씬 어려운 복잡하고 뒤섞인 형태로 코드를 변환하는 보호 기법입니다. 프로그램 본래의 로직과 목적을 숨김으로써 리버스 엔지니어링을 방어하는 것이 주된 목적입니다.