

4단원. 침투 테스트 기초

대상: 전공자 · 정보처리기사 취득자 | 목적: 인프라 직무(네트워크·클라우드·시스템·보안) 기술 면접 대비
중요도: ★★★ 최빈출 · ★★ 자주 출제 · ★ 기본 개념

Q1. 모의해킹(침투 테스트)의 일반적인 절차를 단계별로 설명하세요. ★★★

답안 모의해킹은 일반적으로 정찰(**Reconnaissance**) → 스캐닝(**Scanning**) → 취약점 분석(**Vulnerability Assessment**) → 공격 및 침투(**Exploitation**) → 권한 상승/내부 이동(**Post-Exploitation**) → 보고(**Reporting**)의 순서로 진행됩니다. 정찰 단계에서는 대상의 도메인, IP 대역, 조직 정보 등을 수집하고, 스캐닝 단계에서는 열린 포트와 실행 중인 서비스를 파악합니다. 취약점 분석 단계에서는 발견된 서비스의 알려진 취약점을 확인하고, 공격 단계에서는 사전에 합의된 범위 안에서 실제로 취약점을 악용해 침투 가능성을 검증합니다. 마지막으로 발견한 취약점과 위험도, 개선 방안을 담은 보고서를 작성해 고객에게 전달하는 것으로 마무리되며, 모든 단계는 사전에 서면으로 합의된 범위(Scope)와 규칙(Rules of Engagement) 내에서만 수행되어야 합니다.

관련 개념 정찰/스캐닝/취약점분석/공격/보고, Rules of Engagement, Scope 정의, 화이트박스/블랙박스 테스트, 보고서 작성

Q2. 정찰(**Reconnaissance**) 단계에서 수행하는 활동을 능동적 정찰과 수동적 정찰로 구분해 설명하세요. ★★★

답안 수동적 정찰(**Passive Reconnaissance**)은 대상 시스템에 직접적인 트래픽을 발생시키지 않고 공개된 정보만으로 정보를 수집하는 방식으로, 검색 엔진, WHOIS 조회, DNS 공개 레코드 확인, 소셜미디어·채용공고 분석 등이 해당하며 대상이 이를 탐지하기 어렵다는 특징이 있습니다. 능동적 정찰(**Active Reconnaissance**)은 대상 시스템에 직접 패킷을 보내 응답을 확인하는 방식으로, 포트 스캔이나 서비스 배너 확인 등이 해당하며 방화벽이나 IDS에 탐지될 가능성이 있습니다. 실무 모의해킹에서는 보통 탐지 위험이 낮은 수동적 정찰로 최대한 정보를 확보한 뒤, 필요한 범위 내에서만 능동적 정찰로 넘어가는 순서를 따릅니다.

관련 개념 WHOIS, DNS 열거, OSINT, 포트 스캔, 배너 그레빙

Q3. 취약점 스캐닝(**Vulnerability Scanning**)과 취약점 진단(**Assessment**)의 차이를 설명하세요. ★★★

답안 취약점 스캐닝은 자동화된 도구를 사용해 대상 시스템의 알려진 취약점(CVE 등)을 패턴 매칭 방식으로 빠르게 탐지하는 작업으로, 짧은 시간에 넓은 범위를 점검할 수 있지만 오탐과 미탐이 발생할 수 있습니다. 취약점 진단(평가)은 스캐너의 결과를 바탕으로 실제로 악용 가능한지, 비즈니스에 미치는 영향은 어느 정도인지를 사람이 직접 검증하고 우선순위를 매기는 과정으로, 스캐닝보다 더 정확하고 심층적인 결과를 제공합니다. 실무에서는 정기적인 취약점 스캐닝으로 전체 자산의 상태를 빠르게 파악하고, 위험도가 높다고 판단되는 항목은 별도의 수동 진단이나 침투 테스트로 심층 검증하는 방식으로 운영합니다.

관련 개념 CVE, 오탐/미탐, 자동화 스캐너, 위험도 우선순위(CVSS), 수동 검증

Q4. OWASP Top 10의 개념과 대표적인 취약점 유형 몇 가지를 개념 수준에서 설명하세요. ★★★

답안 **OWASP Top 10**은 웹 애플리케이션에서 발생하는 대표적인 보안 취약점을 주기적으로 정리해 발표하는 커뮤니티 기반 가이드로, 개발자와 보안 담당자가 우선적으로 점검해야 할 취약점 유형을 파악하는 데 널리 활용됩니다. 대표적인 유형으로는 사용자 입력값을 제대로 검증하지 않아 데이터베이스 쿼리가 조작되는 인젝션(**Injection**, 예: **SQL Injection**), 웹 페이지에 악성 스크립트를 삽입해 다른 사용자의 브라우저에서 실행되게 하는 **XSS(Cross-Site Scripting)**, 접근 권한 검증이 미흡해 인가되지 않은 자원에 접근 가능한 취약한 접근 제어(**Broken Access Control**), 그리고 알려진 취약점이 있는 오래된 소프트웨어 컴포넌트를 그대로 사용하는 문제 등이 있습니다. 이 목록은 특정 취약점의 존재 여부보다는 개발·운영 단계에서 놓치기 쉬운 보안 카테고리를 이해하는 데 그 의의가 있습니다.

관련 개념 SQL Injection, XSS, Broken Access Control, 입력값 검증, 시큐어 코딩

Q5. SQL 인젝션 공격의 원리와 방어 방법을 설명하세요. ★★★

답안 **SQL** 인젝션은 애플리케이션이 사용자 입력값을 검증 없이 그대로 SQL 쿼리 문자열에 결합해 실행할 때, 공격자가 입력값에 **SQL** 구문의 일부(따옴표, 논리 연산자 등)를 삽입해 원래 쿼리의 의미를 조작하는 공격입니다. 이를 통해 인증 우회, 비인가 데이터 조회, 데이터 변조·삭제 등이 가능해질 수 있습니다. 근본적인 방어 방법은 사용자 입력을 쿼리 문자열에 직접 이어붙이지 않고 파라미터화된 쿼리(**Prepared Statement**)를 사용해 입력값이 항상 데이터 값으로만 취급되도록 하는 것이며, 추가적으로 입력값 검증(화이트리스트 방식), 데이터베이스 계정의 최소 권한 부여, WAF를 통한 탐지·차단을 병행합니다.

관련 개념 Prepared Statement, 입력값 검증, WAF, 최소 권한 DB 계정, Blind SQL Injection

Q6. XSS(Cross-Site Scripting) 공격의 유형(저장형/반사형/DOM 기반)을 설명하세요. ★★★

답안 저장형(**Stored**) **XSS**는 공격자가 삽입한 악성 스크립트가 게시판 글, 댓글처럼 서버에 영구 저장되어 이후 그 페이지를 열람하는 모든 사용자에게 실행되는 방식으로 피해 범위가 가장 넓습니다. 반사형(**Reflected**) **XSS**는 악성 스크립트가 담긴 URL을 사용자가 클릭하도록 유도해, 서버가 그 요청 값을 그대로 응답 페이지에 반영할 때 즉시 실행되는 방식입니다. **DOM** 기반(**DOM-based**) **XSS**는 서버를 거치지 않고 클라이언트 측 자바스크립트가 사용자 입력을 그대로 **DOM**에 반영할 때 발생합니다. 공통적인 방어 방법은 사용자 입력을 화면에 출력하기 전에 **HTML** 인코딩(이스케이프 처리)을 적용하는 것이며, 추가로 쿠키에 **HttpOnly** 속성을 부여해 스크립트를 통한 쿠키 탈취를 방지하고 **CSP(Content Security Policy)**로 스크립트 실행 범위를 제한합니다.

관련 개념 저장형/반사형/DOM 기반, HTML 인코딩, HttpOnly, CSP, 세션 쿠키 탈취

Q7. 화이트박스, 블랙박스, 그레이박스 테스트의 차이를 설명하세요. ★★

답안 블랙박스(**Black-box**) 테스트는 테스터가 대상 시스템에 대한 사전 정보(소스 코드, 내부 구조 등)를 전혀 제공받지 않고 외부 공격자의 관점에서 진행하는 방식으로, 실제 공격 상황을 가장 유사하게 재현할 수 있지만 시간이 오래 걸립니다. 화이트박스(**White-box**) 테스트는 소스 코드, 아키텍처 문서, 계정 정보 등 내부 정보를 모두 제공받아 진행하는 방식으로, 짧은 시간에 더 깊이 있는 취약점을 찾을 수 있습니다. 그레이박스(**Gray-box**) 테스트는 일부 정보(예: 일반 사용자 계정)만 제공받아 진행하는 절충된 방식으로, 실무에서는 시간과 비용 효율을 고려해 그레이박스 방식이 자주 채택됩니다.

관련 개념 내부 정보 제공 범위, 공격자 관점 시뮬레이션, 테스트 효율성, 소스 코드 리뷰, 계정 기반 테스트

Q8. 윤리적 해킹(Ethical Hacking)의 범위와 법적·윤리적 한계에 대해 설명하세요. ★★

답안 윤리적 해킹은 반드시 대상 조직으로부터 서면으로 명시적인 사전 승인(계약서, **Scope of Work**)을 받은 범위 내에서만 수행되어야 하며, 승인되지 않은 시스템이나 제3자의 자산을 건드리는 것은 정당한 사유 없이는 정보통신망법 등 관련 법률 위반에 해당할 수 있습니다. 테스트 범위(Scope), 허용되는 공격 기법, 테스트 가능 시간대, 비상 연락 체계 등을 사전에 규칙(**Rules of Engagement**)으로 명확히 문서화해야 하며, 테스트 과정에서 발견한 민감 정보는 비밀유지계약(**NDA**)에 따라 외부에 유출하지 않아야 합니다. 또한 서비스 가용성에 실질적 피해를 줄 수 있는 공격(예: 실제 DDoS)은 사전 합의가 없다면 수행해서는 안 되며, 발견된 취약점은 악용하지 않고 즉시 보고해 개선으로 이어지도록 하는 것이 윤리적 해커의 책임입니다.

관련 개념 Rules of Engagement, NDA, 서면 승인, 정보통신망법, 책임있는 취약점 공개(Responsible Disclosure)

Q9. 취약점의 위험도를 평가할 때 사용하는 기준(CVSS 등)의 개념을 설명하세요. ★★

답안 발견된 취약점은 모두 동일하게 다뤄지지 않으며, 얼마나 쉽게 악용될 수 있는지(공격 복잡도, 필요 권한, 사용자 개입 여부)와 악용되었을 때 어떤 피해(기밀성·무결성·가용성 영향)를 주는지를 종합해 위험도를 평가합니다. **CVSS(Common Vulnerability Scoring System)**는 이러한 요소들을 정량화해 0~10점 사이의 점수로 산출하는 국제적으로 널리 사용되는 평가 체계로, 점수 구간에 따라 심각(Critical)·높음(High)·중간(Medium)·낮음(Low) 등으로 분류합니다. 실무에서는 CVSS 점수만이 아니라 해당 시스템이 비즈니스에서 갖는 중요도, 실제 인터넷에 노출되어 있는지 여부 등을 함께 고려해 패치 우선순위를 정합니다.

관련 개념 CVSS 점수 체계, 공격 벡터/복잡도, 심각도 분류, 패치 우선순위, 비즈니스 영향도

Q10. 모의해킹 결과 보고서에는 어떤 내용이 포함되어야 하나요?



답안 보고서는 경영진과 실무 엔지니어 모두가 활용할 수 있도록 구성되어야 합니다. 우선 비전문가도 이해할 수 있도록 전체 결과와 위험 수준을 요약한 경영진 요약 **(Executive Summary)**이 있어야 하고, 이어서 테스트 범위와 방법론을 명시해 어떤 조건 하에서 수행된 결과인지 밝혀야 합니다. 핵심은 발견된 각 취약점에 대해 재현 절차, 위험도(**CVSS** 등), 영향 범위, 구체적인 개선 권고사항을 상세히 기술하는 부분이며, 스크린샷이나 로그와 같은 증적 자료를 함께 첨부해 신뢰성을 높입니다. 마지막으로 전체 취약점을 위험도 순으로 정리한 요약표를 제공해 담당자가 조치 우선순위를 빠르게 파악할 수 있도록 하는 것이 일반적입니다.

관련 개념 Executive Summary, 재현 절차(PoC), 개선 권고사항, 증적 자료, 위험도 순위표
