

위험 체급 기반 동적 사이버보안 실태평가 프레임워크



김홍근 | 동국대학교 국제정보보호대학원

I. 서론

1. 연구의 배경: 데이터 중심 패러다임 전환과 평가 체계 재설계의 필요성
2. 문제 제기: 정적 분류 체계의 한계와 점수 보상 오류
3. 연구 목적 및 기여: 위험 체급 기반 사이버보안 실태평가 프레임워크 제안

II. 이론적 배경 및 관련 연구

1. 현행 실태평가의 점수 합산 구조와 비보상적 모델의 필요성
2. 기존 실태평가 체계 개선 연구와 본 연구의 위치
3. 공격 표면 관리와 3차원 위험 벡터
4. 옐로 레이팅 시스템의 원리와 사이버 위험 평가 적용 타당성

III. 기관 보안 위험 DNA 기반 3차원 위험 상태 모델링

1. 기관 보안 3차원 위험 벡터 및 상태값 정의
2. 기밀 데이터 보유에 대한 하드코딩 트리거 기반 예외 처리 로직

IV. 위험 벡터의 상대 위험 순위화 및 위험 체급 분류 메커니즘

1. 평가 대상의 한정 및 무작위 표본 추출 기반 분산 매치업
2. 옐로 레이팅 알고리즘을 활용한 상대 위험 서열화
3. Fisher-Jenks 알고리즘을 활용한 위험 체급 분류 및 최종 병합
4. 예시 데이터를 활용한 위험 체급 산출 절차

V. 위험 체급별 평가 지표 차등 할당 및 비보상적 평가 체계

1. 위험 체급에 따른 계층형 평가지표 구조 설계
2. 위험 체급별 평가지표 차등 할당 모델
3. 비보상적 평가 및 유연한 강등

VI. 결론 및 향후 연구 과제

1. 연구 결과 요약
2. 연구의 의의와 정책적 함의
3. 연구의 한계점 및 향후 과제

논문 요약

본 연구는 망분리 완화와 다중계층보안(MLS) 기반 N2SF 도입 환경에서 공공기관의 보안 위험 수준에 따라 평가지표의 범위와 강도를 차등화하는 위험 체급 기반 동적 사이버보안 실태평가 프레임워크를 제안한다. 기존 체크리스트 기반 실태평가는 항목별 이행 점수를 합산하여 종합 결과를 산출하는 구조이므로, 핵심 보안 결함이 다른 행정·관리 지표의 높은 점수로 상쇄되는 점수 보상 및 위험 희석 문제가 발생할 수 있다. 또한 모든 기관에 동일한 평가 기준을 적용하는 방식은 기관별 데이터 중요도, 외부 연결성, 방어 역량의 차이를 충분히 반영하기 어렵다. 이를 개선하기 위해 본 연구는 기관의 위험 상태를 데이터 중력(DGI), 망 복잡도(NCI), 방어 역량(ACI)으로 구성된 위험 DNA 벡터로 모델링하였다. 기밀 데이터 보유 등 결정적 위험 요인은 최상위 체급으로 우선 분류하고, 나머지 위험 벡터는 전문가 쌍대 비교와 Elo 레이팅을 통해 1차원 상대 위험 순위로 변환하였다. 이후 Fisher-Jenks 알고리즘을 적용하여 3개 위험 체급을 분류하고, 산출된 체급에 따라 계층형 평가지표를 차등 할당하였다. 또한 임계치 과락과 핵심 지표 미충족 등 비보상적 평가 로직을 적용함으로써 핵심 위험요소가 다른 항목의 양호한 평가 결과로 상쇄되는 문제를 완화하고자 하였다. 본 연구는 기관별 위험 특성에 따라 평가 강도와 평가지표 범위를 사전에 조정하고, 공격 표면 확대에 따른 보안통제 요구를 평가체계에 반영하는 위험 기반 실태평가 프레임워크를 제시한다는 점에서 의의가 있다.

주제어 국가망보안프레임워크, 사이버보안실태평가, 위험 DNA, 엘로 레이팅 알고리즘, 위험 체급

I. 서론

1. 연구의 배경:

데이터 중심 패러다임 전환과 평가 체계 재설계의 필요성

현대 국가 안보의 핵심 기반인 공공 부문의 정보기술 환경은 본격적인 패러다임 전환기를 맞이하고 있다. 그동안 공공분야 사이버보안의 근간을 이루었던 물리적 망분리 정책은 클라우드 컴퓨팅, 생성형 인공지능, 빅데이터 등 신기술 도입이 가속화됨에 따라 그 효용성의 한계를 드러내고 있다. 이에 따라 자산의 물리적 위치가 아니라 데이터의 중요도(비밀·민감·공개)에 기반하여 보안 강도를 차등 적용하는 다중계층보안(Multi-Level Security, MLS) 중심의 차세대 국가망보안 프레임워크(N2SF)로 거버넌스 전환이 추진되고 있다(국가정보원 2025; 한국인터넷진흥원 2026).

이러한 데이터 중심(Data-centric)의 사이버보안 패러다임은 보안의 유연성과 업무 효율성을 높일 수 있지만, 동시에 외부망과의 연결점, 즉 공격 표면(Attack Surface)을 복합적으로 확장시킬 수 있다. 따라서 사이버보안 거버넌스의 관점도 형식적 규제 준수(Compliance) 중심에서 실질적 위험 기반 방어(Risk-based Defense) 중심으로 전환될 필요가 있다. 이는 글로벌 표준인 NIST CSF 2.0(National Institute of Standards and Technology 2024)이 제시하는 위험 관리(Risk Management) 중심의 거버넌스 체계와도 그 방향을 같이한다.

이를 위해서는 기존의 일괄적 평가 방식에서 벗어나, 보유 데이터의 중요도, 네트워크 연결성, 방어 역량 등 기관별 위험 요인을 종합

적으로 고려하여 보안 위험 수준을 분류하고, 해당 위험 수준에 상응하는 평가지표를 차등 적용할 수 있는 차세대 평가 프레임워크가 필요하다(한국사이버안보학회 2026).

2. 문제 제기: 정적 분류 체계의 한계와 점수 보상 오류

성공적인 차세대 보안 거버넌스로의 전환을 위해서는, 먼저 현행 공공 정보보안 실태평가 모델이 내포하고 있는 다음과 같은 구조적 한계를 식별하고 개선할 필요가 있다.

첫째, 일괄적이고 정적인 평가지표(One-Size-Fits-All) 적용의 한계이다. 현행 평가는 개별 기관이 취급하는 데이터의 민감도나 네트워크의 복잡도를 충분히 반영하기보다, 기관 규모 및 행정적 분류를 기준으로 평가 그룹을 구분한다. 이로 인해 고위험 데이터를 취급하는 중소규모 특수 기관의 실질적 위험이 과소평가되는 보안 수준의 착시가 발생할 수 있다.

둘째, 단순 산술 합산 방식이 초래하는 점수 보상(Score Compensation) 오류 및 평균의 함정이다. 현행 실태평가는 관리적, 물리적, 기술적 통제 지표를 개별 평가항목으로 구성하고, 각 항목의 점수를 합산하여 종합 결과를 산출하는 구조를 가진다. 이 구조에서는 클라우드 접근통제 실패나 제로 트러스트 인증 누락과 같은 주요 기술 결함(Critical Defect)이 다른 관리적·행정적 통제 항목의 높은 점수로 상쇄되어, 전체 평가 결과에서 핵심 위험이 충분히 드러나지 않는 문제가 발생할 수 있다. 공격자는 종합점수가 아니라 실제 취약점을 공략한다는 점에서, 이러한 보상적 평가 구조는 실질적 위험 기반 평가와 괴리될 수 있다.

3. 연구 목적 및 기여:

위험 체급 기반 사이버보안 실태평가 프레임워크 제안

본 연구의 목적은 기존 정적 평가 체계가 가진 구조적 한계를 해결하기 위해, 망분리 정책 완화 및 MLS 체계인 N2SF 도입 환경에 맞춰 공공기관의 보안 위험 수준을 위험 체급(Risk Tier)으로 분류하고, 해당 위험 수준에 따라 평가지표를 차등 적용할 수 있는 동적 사이버보안 실태평가 프레임워크를 제안하는 데 있다. 본 연구가 제안하는 프레임워크의 주요 구성 요소와 기여점은 다음과 같다.

- 위험 DNA(Risk DNA) 기반 벡터 모델링: 기관의 고유 특성을 데이터 중력(DGI), 망 복잡도(NCI), 방어 역량(ACI)의 세 가지 관점에서 구조화하여, 단순한 기관 규모나 행정적 분류가 아니라 기관이 처한 위험 상태를 3차원 위험 벡터로 표현한다.
- Elo 레이팅 기반 상대 위험 서열화 및 체급화 절차 적용: 객체 간 상대적 우위를 산출하는 통계적 알고리즘인 Elo 레이팅 시스템(Elo 1978)을 위험 벡터 간 상대 위험 판단에 응용하여, 기관 위험 벡터의 1차원 상대 위험 순위를 도출한다. 이후 정책적·행정적 수용성을 고려해 사전 설정한 3개 위험 체급에 대해 Fisher-Jenks 알고리즘(Jenks 1967)을 적용함으로써, 체급 내 동질성과 체급 간 구분성이 높은 위험군 분류를 수행한다.
- 비보상적(Non-Compensatory) 룰 기반 계층형 평가지표 할당: 산출된 위험 체급에 따라 3 계층(Layer) 기반의 계층형 평가지표를 차등 할당하고, 기밀 데이터 보유 시 최상위 체급으로 우선 분류하는 규칙 및 임계치 과락(Minimum Threshold Failure) 등 비보상적 로직을 적용한다. 이를 통해 특정 핵심 위험요소가 다른 항목

의 양호한 평가 결과로 상쇄되는 문제를 완화하고, 기관의 위험 특성에 부합하는 평가 강도와 지표 범위를 설정할 수 있도록 한다.

다만 본 연구는 체급별로 적용될 개별 평가지표의 세부 목록을 확정하거나, 각 지표의 구체적 측정·판정 방식을 설계하는 것을 직접적인 연구 범위로 삼지는 않는다. 본 연구의 초점은 기관이 처한 위험 수준에 따라 어떤 범위와 강도의 평가지표를 적용할 것인지를 사전에 결정하는 위험 체급화 및 지표 할당 프레임워크를 제안하는 데 있다. 따라서 개별 지표의 점수화 방식, 적합성 판정 기준, 증빙자료 요건 및 감사 절차는 후속 연구와 제도화 단계에서 구체화되어야 한다.

II. 이론적 배경 및 관련 연구

1. 현행 실태평가의 점수 합산 구조와 비보상적 모델의 필요성

현행 정보보안 실태평가 방식은 관리적, 물리적, 기술적 보안 통제 항목을 개별 평가항목으로 구성하고, 각 항목의 점수를 합산하여 종합 결과를 산출하는 구조를 가진다. 이러한 점수 합산 구조는 평가 결과를 산출하고 기관 간 비교를 수행하는 데 행정적 편의성을 제공하지만, 핵심 보안통제의 미흡이 다른 항목의 높은 점수로 상쇄될 수 있다는 한계를 지닌다. 특히 복합적인 클라우드 및 하이브리드 네트워크 환경에서는 개별 취약점의 상호작용(Vulnerability Chaining)으로 인해 위험이 비선형적으로 증폭될 수 있다. 따라서 각 보안통제 항목의

점수를 단순 합산하는 방식은 통제 실패 간 연쇄 효과와 위험 증폭 가능성을 충분히 반영하기 어렵다.

이러한 한계를 고려할 때, 차세대 보안 실태평가 체계는 핵심 보안 통제의 미흡이 다른 평가항목의 높은 점수로 상쇄되는 구조를 완화할 필요가 있다. 이를 위해 특정 핵심 지표의 임계치 미달이나 필수 통제항목의 미흡이 다른 항목의 점수로 보상되지 않도록 하는 비보상적(Non-Compensatory) 평가 로직이 요구된다.

2. 기존 실태평가 체계 개선 연구와 본 연구의 위치

사이버 보안 실태평가 체계의 개선 필요성은 기존 연구에서도 제기되어 왔다. 최명길(최명길 2025)은 국내 공공기관을 대상으로 시행되는 사이버 보안 실태평가 체계가 일률적 체크리스트 항목 중심으로 구성되어 있어 기관별 보안역량이나 위험요인을 충분히 반영하지 못한다고 지적하였다. 해당 연구는 기존 평가체계의 한계를 보완하기 위해 기관 특성과 보안환경을 반영할 수 있는 자율 기반 사이버보안 평가체계를 제안하였다.

최명길은 이러한 문제의식을 뒷받침하기 위해 국내외 보안 평가체계를 비교하였다. 국내 평가체계로는 현행 국정원의 사이버보안 실태평가, ISMS-P, CSAP, 주요정보통신기반시설 이행점검 등을 검토하고, 국외 평가체계로는 미국 FISMA와 FedRAMP, EU NIS2, 영국 GovS 007 및 CAF, 독일 IT-Grundschutz, 일본 JISMS, ISO/IEC 27001 등을 분석하였다. 해당 연구는 특히 ISMS-P, FISMA, FedRAMP 등은 기관 특성에 따라 일정 부분 평가범위나 통제항목을 조정할 수 있는 개방형 평가체계의 성격을 갖는 반면, 현행 사이버보안

실태평가는 국가가 정한 폐쇄형 평가지표를 모든 기관에 공통 적용하는 구조라는 점을 지적하였다. 이러한 비교는 기존 실태평가 체계가 기관별 보안환경과 위험 수준을 충분히 반영하기 어렵다는 본 연구의 문제의식과 연결된다.

다만 최명길의 제안은 새로운 평가 체계라기보다는, ISO/IEC 27001 및 ISMS 계열에서 널리 활용되는 위험평가와 PDCA(Plan-Do-Check-Act) 수명주기 원리를 공공기관 사이버보안 실태평가 맥락에 적용·재구성한 접근으로 이해할 수 있다. 해당 연구는 위험평가를 통해 기관별 보안환경과 위협을 분석하고, 이를 바탕으로 계획 수립, 실행, 평가, 개선의 순환 절차를 수행하는 평가체계의 방향을 제시한다. 이는 기존의 일률적 체크리스트 평가를 기관 특성과 위험관리 중심으로 전환해야 한다는 점에서 중요한 선행연구이다.

그러나 위험평가와 PDCA 기반 접근은 기관이 보안관리 활동을 어떻게 계획·운영·개선할 것인지에 대한 관리체계적 틀을 제공하는 데 강점이 있는 반면, 기관 간 위험 수준을 어떤 방식으로 상대 서열화하고, 그 결과를 평가 지표의 차등 배정으로 어떻게 연결할 것인지에 대해서는 추가적인 구체화가 필요하다. 특히 기존 실태평가에서 제기되는 점수 보상 문제, 고위험 데이터 보유 기관의 과소평가 가능성, 기관별 공격 표면 차이에 따른 평가 강도 조정 문제를 해결하기 위해서는 위험평가 절차만으로는 부족하며, 위험 프로파일의 구조화와 체급화 메커니즘이 필요하다.

본 연구는 이러한 선행연구의 문제의식을 계승하되, 실태평가 2.0 개발 요구사항(한국사이버안보학회 2026)에 대응하여 기관의 위험 수준을 체급화하고, 그 결과를 평가지표의 차등 할당으로 연결하는 절차적 메커니즘을 제시한다. 구체적으로 기관의 위험 특성을 데이터 중

력(DGI), 망 복잡도(NCI), 방어 역량(ACI)의 3차원 위협 DNA로 모델링하고, 전문가 쌍대 비교와 Elo 레이팅을 활용하여 위협 벡터 간 상대 위험 순위를 산출한다. 이후 Fisher-Jenks 알고리즘을 통해 정책적·행정적 수용성을 고려해 사전 설정한 3개 위험 체급으로 분류하고, 체급별로 계층형 보안 평가지표를 차등 할당한다. 또한 C등급 데이터 보유 시 최상위 체급 우선 분류, 임계치 과락, 핵심 지표 기반 평가 결과 상한 제한 등 비보상적 평가 로직을 도입하여 기존 합산식 평가에서 발생하는 점수 보상 문제를 완화한다.

따라서 본 연구는 기존 실태평가 개선 연구가 제기한 기관 특성 반영의 필요성을 출발점으로 삼아, 기관 위험 프로파일을 상태값 기반 위험 벡터로 구조화하고, 이를 상대 서열화·체급화한 뒤, 체급별 평가 지표 배분과 비보상적 평가 로직으로 연결하는 절차적 프레임워크를 제안한다는 데 차별성이 있다.

3. 공격 표면 관리와 3차원 위협 벡터

차세대 N2SF와 같이 데이터 중심의 유연한 망 구성이 확대됨에 따라, 기관의 외부 연결점, 즉 공격 표면(Attack Surface)의 관리가 핵심 과제로 대두되었다. 자산의 경계가 모호해진 환경에서는 단편적인 체크리스트 중심의 지표만으로는 실질적인 방어 태세를 평가하기 어렵다. 동적으로 변화하는 공격 표면을 평가하고 관리하기 위해서는 기관이 직면한 위협을 단일 점수가 아니라 다차원적 위험 상태로 해석할 필요가 있다. 즉, 기관이 보호해야 할 데이터의 민감도와 규모, 클라우드 및 AI 도입으로 인해 확장된 네트워크의 복잡도, 이에 대응하는 관리적 방어 역량을 종합적으로 고려해야 한다. 본 연구는 이 세

가지 요소를 핵심 파라미터로 하는 3차원 위험 벡터(Risk Vector) 개념을 도입한다. 개별 기관의 고유한 보안 환경을 3차원 벡터 공간에 표현함으로써, 기관의 공격 표면과 방어 조건이 결합된 위험 상태를 구조적으로 해석할 수 있는 이론적 기반을 마련한다.

4. 엘로 레이팅 시스템의 원리와 사이버 위험 평가 적용 타당성

엘로 레이팅 시스템은 본래 다수의 평가 대상이 존재하는 경쟁 환경에서 객체 간의 상대적 우위를 확률적으로 산출하기 위해 고안된 통계적 알고리즘이다. 이 시스템은 두 객체의 현재 점수 격차를 바탕으로 승리 기댓값을 로지스틱 곡선(Logistic Curve) 형태로 산출하며, 실제 쌍대 비교 결과와 기댓값의 오차를 반영하여 점수를 지속적으로 갱신하는 원리를 가진다.

본 연구는 이러한 통계적 특성을 활용하여, 전문가 쌍대 비교 결과를 바탕으로 3차원 위험 벡터 간 상대 위험 서열을 산출한다. DGI, NCI, ACI로 구성된 3차원 위험 벡터를 바탕으로 전문가 패널이 위험 벡터 간 상대적 위험도를 쌍대 비교(Pairwise Comparison)하고, 그 판정 결과를 Elo 알고리즘의 입력값으로 활용한다. 알고리즘은 이 비교 결과를 누적하여 위험 벡터의 1차원 상대 위험 서열을 도출한다. 이렇게 도출된 서열 점수는 후속 단계에서 Fisher-Jenks 알고리즘에 입력되어 위험 체급(Tier) 분류에 활용되며, 각 체급에 상응하는 계층형 평가표 차등 할당의 기준으로 사용된다.

이와 같은 접근은 사이버보안 분야에서 전례 없는 임의적 차용이 아니라, CVSS v4.0 점수체계 개발 과정에서 활용된 전문가 쌍대 비교 기반 서열화 방식과 방법론적 유사성을 가진다. FIRST의 CVSS

v4.0 문서(FIRST 2023a; FIRST 2023b)에 따르면, CVSS v4.0은 약 1,500만 개의 CVSS-BTE 벡터를 270개의 동치 집합으로 축약한 뒤, 전문가들이 각 동치 집합을 대표하는 벡터들의 상대적 심각도를 비교하도록 하였고, 이 비교 결과를 바탕으로 벡터 집합을 낮은 심각도에서 높은 심각도로 정렬하였다. 이 과정에서 전문가 비교 결과는 일종의 경기 결과처럼 취급되었으며, 체스 경기 결과로 선수 순위를 산출하는 알고리즘과 동일한 원리의 엘로 알고리즘이 벡터 집합의 심각도 순위 산출에 사용되었다. 따라서 본 연구의 엘로 적용은 게임 레이팅 알고리즘의 단순한 은유적 차용이 아니라, 다차원 보안 속성 벡터를 전문가 판단에 기반한 1차원 순위로 변환한 CVSS v4.0의 방법론적 선례를 기관 위험 평가 문제에 적용한 것이다. 다만 CVSS v4.0이 취약점 벡터의 심각도 순위를 산출하는 데 초점을 둔 반면, 본 연구는 기관의 위험 상태를 표현한 3차원 위험 벡터를 상대 위험 서열로 변환한다는 점에서 적용 대상이 다르다.

III. 기관 보안 위험 DNA 기반 3차원 위험 상태 모델링

1. 기관 보안 3차원 위험 벡터 및 상태값 정의

본 연구에서 제시하는 모델은 기관의 보안 위험 특성을 데이터(Data), 네트워크(Network), 방어 역량(Administrative Capacity)의 세 가지 핵심 관점에서 프로파일링하는 ‘위험 DNA’ 프레임워크를 기반으로

한다. 이는 현대 위험 평가의 기본 구조인 $Risk = Impact \times Likelihood$ (National Institute of Standards and Technology 2012)를 공공기관 사이버보안 실태평가 환경에 맞게 3차원 위험 벡터로 구체화한 것이다.

일반적인 위험평가에서 영향도(Impact)는 사고가 발생했을 때 초래되는 피해의 크기를 의미하고, 발생 가능성(Likelihood)은 사고가 실제로 발생할 가능성을 의미한다. 그러나 공공기관 사이버보안 실태평가에서는 이 두 요소만으로 기관의 위험 상태를 충분히 설명하기 어렵다. 기관이 보유한 데이터의 중요도는 사고 발생 시 피해 규모를 결정하고, 네트워크 구조와 외부 연결성은 공격 표면과 사고 발생 가능성을 변화시키며, 조직의 보안 거버넌스와 방어 역량은 식별된 위험을 완화하거나 통제하는 역할을 수행하기 때문이다.

이에 본 연구는 공공기관의 위험 상태를 세 가지 축으로 구조화한다. 첫째, 데이터 관점은 사고 발생 시 국가안보, 공공서비스, 개인정보 및 민감정보에 미치는 영향도를 반영한다. 둘째, 네트워크 관점은 외부 접점, 클라우드 활용, API 연계, 망간 자료전송 구조 등 공격 표면 확대에 따른 사고 발생 가능성을 반영한다. 셋째, 방어 역량 관점은 기관이 식별된 위험을 통제하고 사고 발생 가능성을 낮출 수 있는 조직적·관리적 완화 능력을 반영한다.

본 연구는 이 세 축을 각각 데이터 중력(DGI: Data Gravity Index), 망 복잡도(NCI: Network Complexity Index), 방어 역량(ACI: Administrative Capacity Index)으로 정의하고, 이들의 조합을 기관 보안 3차원 위험 벡터로 사용한다. DGI와 NCI는 각각 영향도와 발생 가능성의 주요 구성요소를 반영하며, ACI는 위험을 직접 증가시키는 요인이 아니라 위험을 완화하는 제동 변수로 작동한다. 따라서 ACI는 DGI 및 NCI와 위험 벡터의 방향성을 일치시키기 위해 방어 역량이 낮을수록 더

높은 위험 상태값을 부여한다.

이와 같이 정의된 기관의 보안 위험 벡터는 단순한 총점 산식을 구성하기 위한 세부 항목의 나열이 아니라, 기관이 처한 위험 상태를 다차원적으로 표현하기 위한 구조적 프로파일이다. DGI, NCI, ACI의 조합은 기관의 보안 위험을 하나의 점수로 즉시 환산하기 위한 것이 아니라, 전문가 쌍대 비교와 Elo 기반 상대 위험 서열화의 입력 벡터로 활용된다. 즉, 본 연구는 기관별 보안 위험을 “몇 점인가”로 단순 환산하기보다, “어떤 위험 특성 조합을 가진 기관에 더 높은 평가 강도가 필요한가”를 판단하기 위한 위험 DNA 공간을 구성한다.

다만 위험 메트릭(Metric)의 상태값 정의는 완전히 기계적인 산식만으로 결정되기 어렵고, 평가 대상의 맥락과 평가자의 해석을 일정 부분 포함할 수밖에 없다. 대표적으로 CVSS v4.0의 Specification Document도 Attack Vector, Attack Complexity, Privileges Required, User Interaction 등 메트릭과 각 메트릭 값을 정의하지만, 실제 취약점에 어떤 값을 부여할지는 취약점의 공격 조건, 시스템 환경, 영향 범위에 대한 해석을 필요로 한다. 이러한 추상성을 보완하기 위해 CVSS는 Specification Document와 별도로 User Guide, Examples, Calculator 등을 제공하여 사용자가 메트릭 값을 일관되게 판단할 수 있도록 지원한다. 이와 유사하게 본 연구의 DGI·NCI·ACI도 본 논문에서는 위험 DNA 메트릭의 기본 정의와 상태값 분류 기준을 제시하되, 실제 제도화 단계에서는 각 메트릭 값을 구체적 사례에 적용하기 위한 세부 설명과 예시를 제공하는 별도의 “위험 DNA 메트릭 평가 가이드”가 필요하다. 이러한 가이드는 각 상태값의 판정 기준, 대표 사례, 경계 사례 발생 시 참고 기준, 예외 처리 기준, 필요 증빙자료,

필요시 평가자 간 판단 일관성 검토 절차 등을 포함해야 하며, 이는 본 연구의 후속 과제로 둔다.

본 연구에서 DGI, NCI, ACI는 기관의 보안 상태를 임의적으로 수치화하기 위한 추상 지표가 아니라, 실태평가 과정에서 확인 가능한 자료를 바탕으로 위험 상태를 범주화하기 위한 조작적 변수이다. <표 3-1>은 DGI·NCI·ACI의 상태값을 판단하기 위한 기본 기준과 활용 가능한 관찰자료를 요약한 것이다.

<표 3-1> DGI·NCI·ACI 상태값 판정을 위한 기준과 활용 가능한 관찰자료 요약

지표	의미	상태값	기본 판단 기준	활용 가능한 관찰자료
DGI	사고 발생 시 영향도	C / H / M / L	보유 데이터의 분류 등급, 민감정보·개인정보 보유 규모, 국가안보 또는 핵심 업무 영향도	데이터 등급분류 결과, 정보자산 목록, 개인정보·민감정보 보유 현황, 업무 중요도 분석
NCI	사고 발생 가능성에 영향을 주는 공격 표면	H / M / L / N	외부망 접점, 클라우드·SaaS 사용, 생성형 AI API 연계, 망간 자료전송 구조, 외부기관 연계 수준	네트워크 구성도, 클라우드 이용 현황, API 연계 목록, 망간 자료전송 정책, 외부 접속 경로
ACI	위험을 완화 할 수 있는 조직적·관리 적 방어 역량	H / M / N	전담 보안 인력, 독립적 보안 예산, CISO의 보안성 검토 및 보류·거부 권한, 보안정책 집행력	보안조직 현황, 인력 배치표, 예산 편성 자료, 보안성 검토 절차, CISO 권한 규정, 내부 감사 결과

1) 데이터 중력(DGI: Data Gravity Index)

데이터 중력은 사이버 보안 사고 발생 시 국가안보 및 사회적 영향, 즉 위험의 영향도를 결정하는 핵심 변수이다. N2SF의 데이터 중심 보안(Data-Centric Security) 철학을 반영하여, 기관이 보유·처리하는 데이터의 분류등급(C/S/O)과 규모에 따라 다음과 같이 4단계로 분류한

다. 특히 기밀(C) 등급 데이터의 보유 여부는 다른 환경적 변수와 독립적으로 최상위 위험도를 결정하는 우선 분류 기준으로 작용한다.

- C(Critical): 국가 기밀 또는 안보 핵심 데이터를 1건이라도 취급하는 상태
- H(High): 대규모(예: 100만 건 이상) 민감 정보(S) 또는 개인정보를 상시 취급하는 상태
- M(Medium): 중소규모 민감 정보 또는 업무용 비공개 데이터를 취급하는 일반 상태
- L(Low): 공개 정보(O) 또는 낮은 민감도의 데이터 위주로 운영되는 상태

DGI의 핵심은 데이터의 양만이 아니라 데이터 사고가 초래할 수 있는 영향도의 크기를 반영하는 데 있다. 따라서 C등급 데이터는 보유 규모가 작더라도 국가안보 또는 핵심 업무에 미치는 영향이 중대할 수 있으므로, 별도의 최상위 체급 우선 분류 기준으로 작동한다. 반면 H, M, L 상태값은 민감정보·개인정보의 규모, 업무 중요도, 데이터 활용 범위 등을 종합하여 판단한다.

2) 망 복잡도(NCI: Network Complexity Index)

망 복잡도는 외부망과의 연결 점점, 즉 공격 표면의 확대 정도를 반영하여 사이버 보안 사고 발생 가능성의 선행 지표로 활용되는 환경 변수이다. 단순히 물리적 망분리 여부를 확인하는 기존 방식을 넘어, 클라우드 네이티브 환경의 SaaS 도입, 생성형 AI API 연동, 망간 자료 전송(CDS) 정책의 복잡도 등을 종합적으로 고려하여 4단계로 분류한다.

- H(High): 생성형 AI API 연동, 멀티 클라우드 활용, 다수의 외부

망 연계 등 가시성 확보가 어려운 환경

- M(Medium): 단일 클라우드(SaaS/PaaS) 도입 또는 제한적인 외부 API 연동이 존재하는 환경
- L(Low): 전통적인 망 연계 시스템(CDS 등) 중심이며 외부 접점이 엄격히 통제되는 환경
- N(None): 외부 접점이 전혀 없는 완전 폐쇄망 환경

NCI는 기관의 네트워크가 얼마나 복잡한가를 단순히 설비 수나 시스템 수로 측정하려는 지표가 아니다. 본 연구에서 NCI는 외부 연결성, 클라우드 및 API 활용, 망간 자료전송 경로, 외부기관 연계 등 공격자가 접근하거나 우회할 수 있는 경로가 얼마나 넓고 복합적인지를 반영한다. 따라서 NCI 상태값은 네트워크 구성도, 외부 연계 목록, 클라우드 이용 현황, API 연동 현황, 망간 자료전송 정책 등을 통해 판단될 수 있다.

3) 방어 역량(ACI: Administrative Capacity Index)

방어 역량은 기관이 식별된 위협을 통제하고, 사고 발생 가능성을 낮추며, 보안 위협을 실질적으로 완화할 수 있는 조직적·관리적 대응능력을 의미한다. 본 연구에서는 전담 보안 인력의 확보 수준, 독립적 보안 예산 편성권, 신규 시스템 도입 시 CISO의 보안성 검토 및 보류·거부 권한, 보안 정책의 집행력 등 실질적인 거버넌스 권한과 운영 역량을 기준으로 방어 역량을 판단한다. ACI는 위협을 증가시키는 요인이 아니라 위협을 완화하는 제동 변수이므로, 위협 벡터의 방향성을 DGI 및 NCI와 일치시키기 위해 방어 역량이 낮을수록 더 높은 위협 상태값을 부여한다.

- H(High Risk): 전담 보안 인력 확보 수준이 법적·정책적 요구 수

준에 미달하고, CISO의 독립적 보안성 검토 권한이나 예산 통제력이 부족하여 보안 통제가 조직 내에서 실질적으로 작동하기 어려운 상태

- M(Medium Risk): 기본적인 보안 인력과 관리 체계는 갖추고 있으나, 일부 인력 부족, 예산 제약, 정보화 부서 종속성 등으로 인해 보안 거버넌스의 독립성과 집행력이 제한되는 상태
- N(Normal Risk): 충분한 전담 보안 인력과 독립적 예산권을 확보하고 있으며, CISO의 보안성 검토 권한과 보안 정책 집행력이 제도적으로 보장되는 표준 상태

ACI는 단순히 보안 인력 수만을 측정하는 지표가 아니라, 기관 내부에서 보안 통제가 실제로 작동할 수 있는 제도적 권한과 집행력을 함께 반영한다. 예를 들어 전담 보안 조직이 존재하더라도 예산 편성권이 없거나, 신규 시스템 도입 시 CISO의 보안성 검토 의견이 실질적으로 반영되지 않는다면 방어 역량은 제한적으로 평가될 수 있다. 따라서 ACI 상태값은 보안조직 현황, 인력 배치, 예산 편성 자료, CISO 권한 규정, 보안성 검토 절차, 내부 감사 결과 등을 종합하여 판단한다.

2. 기밀 데이터 보유에 대한 하드코딩 트리거 기반 예외 처리 로직

단순 산술 합산 시 발생하는 점수 보상 오류를 구조적으로 방지하기 위해 본 연구는 기밀 데이터 보유 여부를 최상위 위험 체급 분류에 우선 반영하는 하드코딩 트리거(Hard-coded Trigger)를 도입한다. 이는 기관의 망 복잡도가 낮거나 방어 역량이 우수하더라도, N2SF

기준 C등급(기밀) 데이터를 1건이라도 취급하는 경우, 전문가 쌍대 비교 및 엘로 알고리즘 기반 상대 서열화 절차를 거치지 않고 해당 기관의 위험 벡터를 최상위 위험 체급(Tier 1)으로 우선 분류하는 논리 구조이다. 이를 통해 국가안보와 핵심 공공업무에 중대한 영향을 미칠 수 있는 기밀 데이터 보유 기관이 다른 환경적 요인이나 관리적 통제 수준에 의해 낮은 위험 체급으로 분류되는 문제를 방지한다.

IV. 위험 벡터의 상대 위험 순위화 및 위험 체급 분류 메커니즘

3장에서 정의된 DGI·NCI·ACI 기반 3차원 위험 벡터는 단순한 정량 점수로 즉시 환산되는 것이 아니라, 전문가 패널의 쌍대 비교 결과를 Elo 알고리즘에 입력함으로써 전체 대상에 대한 상대적 위험 순위로 변환된다. 본 장에서는 전문가의 상대 위험 판단을 바탕으로 통계적 서열화 메커니즘을 적용하고, 이를 위험 체급(Risk Tier) 분류로 연결하는 과정을 제시한다.

1. 평가 대상의 한정 및 무작위 표본 추출 기반 분산 매치업

본 모델에서 정의한 3차원 위험 벡터의 가능한 조합은 총 48개($4 \times 4 \times 3$)이다. 이 중 기밀 데이터 보유에 대한 하드코딩 트리거를 적용받는 12개의 최상위 위험 벡터(DGI: C 보유)는 알고리즘 연산 이전에 Tier 1로 우선 분류된다. 따라서 나머지 36개의 비-C등급 위험 벡터가 본 장에서 다루는 Elo 기반 상대 서열화의 연산 대상이 된다. 이

36개 위험 벡터를 1차원의 상대 위험 서열 점수로 변환하기 위해 전문가 패널이 쌍대 비교(Pairwise Comparison)를 수행할 경우, 전수 비교에 필요한 조합의 수는 총 630회($\approx C_2$)에 달한다.

계층화 분석법(Analytic Hierarchy Process)의 일관성 비율(Consistency Ratio) 연구(Saaty 1980) 및 이산선택실험(Discrete Choice Experiment) 이론(Louviere 외 2000)은 복합 변수를 다루는 반복적 비교 과업에서 평가자의 인지적 부담이 증가할수록 판단 일관성이 저하될 수 있음을 시사한다. 이러한 평가 부담을 완화하기 위해, 본 모델은 전체 비교 조합 중 약 30~40%의 매치업을 무작위로 추출한다. 추출된 매치업 리스트는 전문가 패널에게 1인당 20~30회 수준으로 분산 할당된다. 이를 통해 개별 평가자의 비교 부담을 관리하면서도, 위험 벡터 간 상대 위험 서열을 산출하는 데 필요한 쌍대 비교 입력을 확보한다.

2. 엘로 레이팅 알고리즘을 활용한 상대 위험 서열화

전문가 패널에 의해 수행된 매치업의 쌍대 비교 판정 결과는 본 모델의 핵심인 엘로 레이팅 알고리즘을 통해 각 위험 벡터의 상대 위험 서열 점수를 갱신하는 데 사용된다. 모든 벡터는 초기 도입 시 동일한 기준 점수(예: $R_0 = 1500$)를 부여받는다.

1) 기대 판정값(Expected Value) 산출

엘로 레이팅 알고리즘은 매칭된 두 벡터 A와 B의 현재 점수 R_A , R_B 를 바탕으로, 전문가가 벡터 A를 더 고위험으로 판정할 기댓값 E_A 를 산출한다.

$$E_A = \frac{1}{1 + 10^{(R_B - R_A)/D}}$$

2) 서열 점수의 갱신

전문가의 실제 판정 결과(*Actual*)를 입력받아, 알고리즘은 기댓값과의 오차를 반영하여 새로운 서열 점수(R_{new})를 갱신한다.

$$R_{new} = R_{old} + K \cdot (Actual - Expected)$$

본 알고리즘의 핵심은 개별 쌍대 비교 결과를 순차적으로 누적하여 전체 위험 벡터의 상대 서열을 형성하는 데 있다. 무작위로 추출된 일부 매치업의 판정 결과는 K 값에 따라 각 위험 벡터의 점수에 반영되며, 이를 통해 직접 비교되지 않은 벡터 간에도 공통 비교 대상과의 연결 관계를 통해 상대적 위치가 간접적으로 조정될 수 있다. 전문가의 평가 데이터가 누적될수록 36개 벡터 전체의 점수는 쌍대 비교 결과를 반영한 상대 위험 서열로 점진적으로 정렬된다. 결론적으로 이 과정은 전문가의 정성적 판단을 반복 가능한 상대 서열 점수로 변환하여, 차후 위험 체급 분류 및 평가지표 차등 할당을 위한 1차원 상대 위험 순위를 도출하는 기제로 작용한다.

3) 파라미터 설정

R_0 , D , K 는 전문가 쌍대 비교 결과를 상대 위험 서열 점수로 변환하는 과정에서 점수의 초기값, 기대확률의 민감도, 업데이트 폭을 조절하는 알고리즘 운영 파라미터이다. 먼저 R_0 는 모든 위험 벡터에 동일하게 부여되는 초기 기준점수이다. 본 연구에서는 특정 위험 벡터에 사전 우위를 부여하지 않기 위해 모든 서열화 대상 벡터의 초기 점수를 동일하게 설정한다. R_0 의 절대값 자체는 최종 위험 체급을 의미하지 않으며, 모든 벡터가 동일한 출발점에서 전문가 비교 결과에 따라 상대적으로 분화되도록 하기 위한 기준점으로 기능한다.

D 는 두 벡터 간 점수 차이가 기대 판정값에 미치는 민감도를 조절

하는 스케일링 상수이다. 일반적인 Elo 레이팅에서는 $D=400$ 이 널리 사용되며, 이는 400점의 점수 차이가 약 10:1 수준의 기대 우위로 해석될 수 있도록 하는 기준값이다. D 값은 기댓값 곡선의 기울기를 조절하므로, 실제 전문가 비교 데이터가 축적된 이후에는 D 값 변화에 따른 상대 서열 민감도를 검토할 필요가 있다.

K 는 한 번의 쌍대 비교 결과가 기존 서열 점수에 반영되는 정도를 조절하는 업데이트 계수이다. K 값이 지나치게 크면 개별 비교 결과에 따라 점수가 과도하게 요동할 수 있고, 반대로 지나치게 작으면 제한된 비교 데이터에서 서열 형성 속도가 느려질 수 있다. K 값 역시 실제 전문가 패널 데이터의 규모, 판단 일관성 등에 따라 조정될 수 있는 운영 파라미터이다.

3. Fisher-Jenks 알고리즘을 활용한 위험 체급 분류 및 최종 병합

1) 자연적 단절 기반의 등급 분류

엘로 알고리즘 연산을 통해 36개의 서열화 대상 벡터는 각각 고유한 상대 위험 서열 점수(Elo Rating Score)를 획득하게 된다. 본 연구는 이 점수 분포를 정책적으로 수용 가능하고 실무 관리가 용이한 3개의 위험 체급(Tier 1, Tier 2, Tier 3)으로 분류하기 위해 Fisher-Jenks 알고리즘을 적용한다. 여기서 목표 클래스의 수($k=3$)는 알고리즘이 자동으로 산출한 최적 군집 수가 아니라, 실태평가 제도의 행정적 직관성, 평가 운영의 단순성, 피평가기관의 정책적 수용성을 고려하여 사전에 설정한 정책 파라미터이다. 즉, 본 연구에서 Fisher-Jenks 알고리즘의 역할은 “몇 개의 체급으로 나눌 것인가”를 결정하는 것이 아니라, 사전 지정된 3개 체급을 전제로 각 체급의 경계값을 점수 분포에

맞게 산출하는 데 있다.

Fisher-Jenks 알고리즘은 1차원 연속형 점수 분포에서 체급 내 분산은 최소화하고 체급 간 분산은 최대화하는 자연적 단절점(Natural Breaks)을 탐색한다. 이는 단순한 등간격 분할이나 임의적 점수 컷오프(Cut-off) 방식과 달리, 실제 점수 분포상 벡터들이 밀집된 구간과 상대적으로 희소한 구간을 반영하여 체급 경계를 설정하는 방식이다. 따라서 유사한 상대 위험 수준을 보이는 벡터들은 동일 체급 내에 배치되고, 위험 특성의 차이가 상대적으로 커지는 구간에서 체급 경계가 형성된다. 이러한 분류 방식은 미세한 점수 차이만으로 체급이 갈리는 절벽 효과(Cliff Effect)를 완전히 제거하는 것은 아니지만, 임의적 기준선에 따른 등급 분할보다 점수 분포의 구조를 더 잘 반영함으로써 체급 분류의 설명 가능성과 정책적 수용성을 높이는 데 기여한다.

2) 최종 위험 체급 확정

Fisher-Jenks 알고리즘을 통해 36개의 상대 서열화 대상 벡터는 Tier 1, Tier 2, Tier 3의 세 위험 체급으로 분류된다. 이때 도출된 체급은 옐로 점수 자체를 최종 보안 점수로 사용하는 것이 아니라, 상대 위험 순위를 기반으로 평가 강도의 차등 적용 범위를 정하기 위한 정책적 분류 결과이다. 이후 3장에서 하드코딩 트리거로 옐로 연산 대상에서 제외되었던 12개의 최상위 위험 벡터, 즉 DGI가 C등급인 벡터들을 Tier 1 그룹에 최종 병합한다. 이는 기밀 데이터 또는 국가 안보상 핵심 데이터를 보유한 기관이 망 복잡도나 방어 역량과 같은 다른 요인에 의해 낮은 체급으로 분류되는 것을 방지하기 위한 최상위 체급 우선 분류 규칙이다.

결과적으로 전체 48개 위험 DNA 조합은 두 단계의 절차를 거쳐

최종 위험 체급으로 확정된다. 첫째, C등급 데이터 보유 벡터는 기밀 데이터 보유 트리거에 의해 Tier 1로 분류된다. 둘째, 나머지 36개 벡터는 전문가 쌍대 비교, 엘로 기반 상대 위험 순위화, Fisher-Jenks 기반 자연적 단절 분류를 거쳐 Tier 1, Tier 2, Tier 3 중 하나로 배정된다. 이와 같이 확정된 위험 체급은 기관의 우열을 단순히 점수화하기 위한 목적이 아니라, 이후 계층형 보안 평가지표를 차등 할당하기 위한 기준으로 활용된다. 즉, 최종 위험 체급은 기관의 데이터 중요도, 망 복잡도, 방어 역량을 종합적으로 반영한 평가 강도 결정 기준이며, 제5장에서 제시하는 Layer 1, Layer 2, Layer 3 평가지표 할당의 직접적인 입력값으로 기능한다.

4. 예시 데이터를 활용한 위험 체급 산출 절차

본 연구는 위험 체급 기반 실태평가 프레임워크를 제안하는 것을 목적으로 하므로, 실제 전문가 패널의 쌍대 비교 데이터를 활용한 실증 분석은 수행하지 않는다. 다만 제안 프레임워크의 적용 절차를 구체적으로 설명하기 위해, 본 절에서는 DGI, NCI, ACI 상태값의 조합으로 구성된 위험 벡터를 사용한 예시 산출 과정을 제시한다. 즉, 본 절의 산출 과정은 [위험 벡터 대상 정의] → [가상의 전문가 쌍대 비교 판정 데이터 생성] → [Elo 레이팅 기반 상대 위험 서열 점수 산출] → [Fisher-Jenks 알고리즘 기반 위험 체급 분할]의 네 단계 절차로 구성된다.

1) 위험 벡터 대상 정의

전체 48개 위험 벡터 중에서, 최상위 위험군(Tier 1)으로 사전 분류

된 12개 벡터를 제외한 나머지 36개의 위험 벡터를 대상으로 Elo 레이팅 기반 서열화 연산이 수행된다. <표 4-1>은 제3장의 정의에 기반하여 본 예시 산출 절차에서 사용되는 위험 벡터의 구성과 엘로 연산 대상을 요약하여 보여준다.

<표 4-1> 48개 위험 벡터 중 일부 예시

벡터 ID	DGI	NCI	ACI	비고
V01	C	H	H	하드코딩 트리거 → Tier 1
V02	C	H	M	하드코딩 트리거 → Tier 1
...				
V13	H	H	H	Elo 레이팅 기반 서열화 대상
V14	H	H	M	Elo 레이팅 기반 서열화 대상
...				
V47	L	N	M	Elo 레이팅 기반 서열화 대상
V48	L	N	N	Elo 레이팅 기반 서열화 대상

2) 가상의 전문가 쌍대 비교 판정 데이터 생성

본 단계에서는 제4장 1절의 평가 설계에 따라 위험 벡터 간 무작위 추출 기반 매치업 데이터셋을 생성하였다. 전체 가능한 비교 조합은 630개이며, 본 예시에서는 이 중 약 35%에 해당하는 220개의 매치업을 추출하였다. 전문가 1인당 비교 부담을 완화하기 위해 할당 매치업 수는 22회로 제한하였으며, 본 예시에서는 10인의 가상 전문가 패널이 쌍대 비교를 수행한 것으로 가정하여 판정 데이터를 구성하였다. 이 판정 데이터는 실제 전문가 조사 결과가 아니라, 제안 프레임워크의 입력·처리·출력 절차를 설명하기 위한 예시 데이터이다. 이와 같이 구성된 220개의 쌍대 비교 판정 데이터는 후속 엘로 기반

상대 위험 서열화의 입력값으로 활용된다. <표 4-2>는 본 절에서 구성한 쌍대 비교 판정 데이터셋의 일부를 제시한 것이다.

<표 4-2> 가상 전문가 쌍대 비교 예시

Match ID	전문가	벡터 A	벡터 B	전문가 판정(A Win=1, Loss=0)
M001	전문가 1	H-H-H	H-M-N	1
M002	전문가 1	M-H-H	H-L-N	1
...	...	...	...	...
M022	전문가 1	H-H-M	M-H-H	1
M023	전문가 2	H-M-H	M-M-M	1
...	...	...	...	...
M103	전문가 5	M-M-L	H-H-H	0
...	...	...	...	...
M219	전문가 10	H-M-H	H-M-N	1
M220	전문가 10	M-M-M	M-H-M	0

3) Elo 레이팅 기반 상대 위험 서열 점수 산출

본 단계에서는 표 4-2의 쌍대 비교 판정 데이터셋을 엘로 레이팅 알고리즘의 입력값으로 사용하여, 서열화 대상인 36개 위험 벡터의 상대 위험 서열 점수를 산출한다. 여기서 엘로 점수는 기관의 절대적 보안점수나 최종 평가점수가 아니라, 쌍대 비교 판정 결과를 누적하여 위험 벡터 간 상대적 위험 우위를 표현하기 위한 중간 서열화 점수이다. 각 매치업에서 더 위험한 것으로 판정된 위험 벡터는 기대 판정값과 실제 판정 결과의 차이에 따라 점수가 증가하고, 상대적으로 낮은 위험으로 판정된 벡터는 점수가 감소한다. 이러한 갱신 과정을 220회의 매치업에 순차적으로 적용함으로써, 36개 위험 벡터는 쌍대

비교 판정 결과에 기반한 1차원 상대 위험 순위로 정렬된다.

이렇게 산출된 Elo 점수는 위험 체급을 직접 확정하는 값이 아니라, 다음 단계에서 Fisher-Jenks 알고리즘을 적용하여 Tier 1, Tier 2, Tier 3의 자연적 단절 기반 체급으로 분류하기 위한 입력값으로 활용된다. <표 4-3>은 엘로 레이팅 적용 후 산출된 36개 위험 벡터 중 일부의 최종 엘로 점수와 상대 위험 순위를 제시한 것이다.

<표 4-3> Elo 레이팅 기반 위험 벡터별 상대 위험 서열 점수(일부)

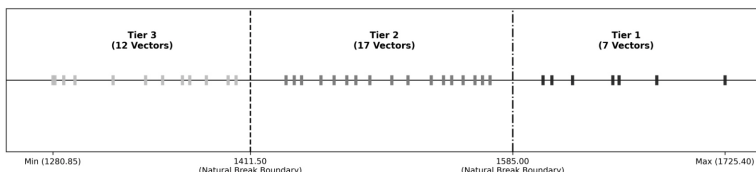
벡터 ID	DGI	NCI	ACI	최종 Elo 점수	상대 위험 순위
V13	H	H	H	1,725.40	1
V14	H	H	M	1,680.15	2
...	...	...	...	...	...
V27	M	H	N	1,505.00	18
...	...	...	...	...	...
V48	L	L	N	1,280.85	36

4) Fisher-Jenks 알고리즘 기반 위험 체급 분할

본 단계에서는 표 4-3에 제시된 36개 위험 벡터의 Elo 점수를 기반으로, Fisher-Jenks 알고리즘을 통해 위험 벡터를 3개의 그룹으로 분할한다. 본 예시에서는 두 개의 체급 경계값으로 1,411.50점과 1,585.00점을 사용하였다. 이에 따라 1,585.00점 이상은 Tier 1, 1,411.50점 이상 1,585.00점 미만은 Tier 2, 1,411.50점 미만은 Tier 3으로 분류된다. 이 경우 36개 위험 벡터는 Tier 1에 7개, Tier 2에 17개, Tier 3에 12개로 분할된다. <그림 4-1>은 이러한 예시 체급 분할 결과를 시각화한 것이다. 그림은 위험 벡터별 엘로 점수의 분포와 Tier 1,

Tier 2, Tier 3의 구분을 보여주며, 점선은 본 예시에서 사용한 Fisher-Jenks 기반 체급 경계를 나타낸다.

〈그림 4-1〉 Fisher-Jenks 알고리즘 기반 위험 체급 분할 예시



V. 위험 체급별 평가 지표 차등 할당 및 비보상적 평가 체계

4장에서 분류된 위험 체급은 단순히 기관의 순위를 매기는 것에 그치지 않고, 각 기관에 적용할 평가지표의 범위와 강도를 결정하는 기준이 된다. 본 장에서는 위험 체급별 계층적 지표 할당 모델과 점수 보상 문제를 완화하기 위한 비보상적 평가 메커니즘을 제시한다.

1. 위험 체급에 따른 계층형 평가지표 구조 설계

본 연구는 모든 공공기관에 동일한 보안 요구사항을 부여하던 기존의 일괄적 평가 방식이 급변하는 정보기술 환경과 기관별 위험 격차를 반영하지 못한다는 한계를 극복하고자 한다. 이를 위해 보안 통제 목적성, 기술적 난이도, 위험 대응의 시급성 등에 따라 평가 지표를 3개의 계층(Layer)으로 구조화한다. 이러한 계층화는 기관이 처

한 위험 체급에 따라 평가지표의 범위와 강도를 차등 적용함으로써, 불필요한 행정 소요를 줄이고 핵심 위험에 방어 역량을 집중시키는 것을 목적으로 한다.

1) Layer 1: 사이버 위생(Cyber Hygiene) 및 기초 통제 계층

모든 공공기관이 예외 없이 준수해야 하는 기본 보안통제 계층이다. 자산 식별, 계정 관리, 패치 관리 등 필수적인 보안 활동들로 구성된다. 글로벌 보안 표준 프레임워크인 CIS(Center for Internet Security)의 Community Defense Model(CDM v2.0)에 따르면(Center for Internet Security 2021), 기본적인 사이버 위생 수준인 IG1(Implementation Group 1)의 이행만으로도 주요 5대 공격 유형에 사용되는 하위 공격 기법(ATT&CK sub-techniques)의 약 77%를 효과적으로 방어할 수 있으며, 전체 보안 통제 구현 시 그 예방률은 91%까지 상승하는 것으로 분석되었다. 이 계층의 지표들은 측정의 편의성을 고려하여 자동화된 로그 확인이나 구성 설정 점검 등 운영상태 확인이 비교적 용이한 항목을 중심으로 배치된다. 이는 모든 보안 체계의 기본 기반이자 상위 계층으로 나아가기 위한 필수 전제 조건이다.

2) Layer 2: 동적 경계 및 컨텍스트 기반 심화 통제 계층

클라우드 전환, 원격 근무 확산 등 경계선이 모호해진 현대적 네트워크 환경에서의 위험 관리에 집중한다. 단순히 “솔루션 도입 유무”를 묻는 수준을 넘어, 신원 기반의 제로 트러스트 인증이나 망간 자료 전송 시의 무해화(CDR, Content Disarm & Reconstruction) 처리 등 실제 데이터의 흐름과 맥락을 통제하는 지표들로 구성된다. 이 계층은 중간 이상 위험 체급(Tier 1, Tier 2) 기관이 직면할 수 있는 확장된 공격

표면과 외부 연계 위험을 평가하기 위한 심화 통제 계층으로 기능한다.

3) Layer 3: 최상위 위험 체급 대상 고도 보안통제 및 신기술 보안 계층

Layer 3는 국가 핵심 자산 또는 고위험 데이터를 보유한 최상위 위험군(Tier 1)을 대상으로 적용되는 고도 보안통제 영역이다. 이 계층은 일반적인 사이버 위생이나 기본 접근통제 수준을 넘어, 고도화된 위협 주체가 활용할 수 있는 우회 경로, 내부 확산 경로, 신기술 기반 공격 표면을 추가적으로 점검하는 데 초점을 둔다. 주요 통제 항목은 AI 모델 무결성 검증, 프롬프트 인젝션 방어, OT/ICS망의 물리적·논리적 가시성 확보, 마이크로 세그멘테이션을 통한 내부 확산 통제, 레드 팀 테스트 및 공격 시뮬레이션 기반 검증 등으로 구성될 수 있다.

이 계층의 목적은 고위험 기관에 대해 강화된 보안통제의 적용 여부와 그 운영 적정성을 평가하는 데 있다. 따라서 Layer 3 지표는 단순한 솔루션 도입 여부를 확인하는 방식보다, 통제 적용 범위, 탐지·차단 효과, 검증 절차의 수행 여부, 개선 조치의 추적 가능성 등 실제 운영 근거를 확인할 수 있는 항목을 중심으로 구성된다. 또한 Layer 3는 Layer 1의 기본 위생 통제와 Layer 2의 동적 경계·컨텍스트 기반 통제가 일정 수준 이상 충족된다는 전제하에 적용되며, 최상위 위험 체급 기관의 평가 강도와 지표 범위를 차등화하기 위한 상위 평가 계층으로 기능한다.

〈표 5-1〉¹⁾은 위험 체급에 따라 차등 적용할 수 있는 계층형 보안 평

1) 본 표는 위험 체급에 따라 적용 가능한 평가지표의 범위와 강도를 예시적으로 제시한 것이다.

가지표의 구성 예시를 제시한다. 이는 개별 지표의 최종 채점 방식이나 감사 절차를 확정하기 위한 것이 아니라, 기관의 위험 수준에 따라 평가지표의 범위와 강도를 차등화할 수 있음을 설명하기 위한 것이다. 본 연구는 이러한 지표 구성을 위해 다음과 같은 세 가지 설계 원칙을 적용하였다..

- **수직적 확장성:** Layer 1은 모든 기관에 공통적으로 요구되는 기본 보안 통제를 평가하고, Layer 2와 Layer 3는 기관의 위험 체급이 높아질수록 요구되는 심화 통제와 고도 보안통제를 평가한다. 이는 평가 대상 기관의 위험 수준에 따라 평가지표의 범위와 강도를 단계적으로 차등화하기 위한 것이다.
- **측정 중심의 지표 구성:** 각 계층의 지표 예시는 단순히 정책 보유 여부를 확인하는 방식에 그치지 않고, 식별률, 적용률, 탐지율, 조치율 등 운영 상태를 확인할 수 있는 측정 가능 지표를 중심으로 구성한다.
- **적층형 방어(Layered Defense) 구조:** 상위 계층의 통제는 하위 계층의 보안통제가 일정 수준 이상 충족된다는 전제하에 의미를 가진다. 따라서 상위 체급 기관일수록 하위 계층의 핵심 통제 미흡이 전체 평가에 미치는 영향을 크게 반영해야 하며, 이는 비보상적 평가 로직의 근거가 된다.

개별 지표의 세부 측정 방식, 적합성 판정 기준, 증빙자료 요건 및 감사 절차는 별도의 평가 지침 또는 후속 연구를 통해 구체화될 필요가 있다.

〈표 5-1〉 체급별 평가 지표 할당 예시

계층	적용 대상	주요 확인 관점	대표 평가 지표 예시
Layer 1: 사이버 위생 및 기초 통제 계층	모든 기관 공통 적용	기본 보안통제가 제도적으로 수립되어 있고 실제 운영 증적이 존재하는지 확인	자산 식별 및 목록 관리, 계정·권한 관리, 패치 및 취약점 관리, 백업 및 복구체계, 로그 수집 및 보관, 침해사고 대응 절차
Layer 2: 동적 경계 및 컨텍스트 기반 심화 통제 계층	중간 이상 위험 체급 기관	기관의 연결성 및 공격 표면 확대 요인에 대해 추가 통제가 적용되고 있는지 확인	외부 연계 시스템 관리, 클라우드·SaaS 보안 설정, API 접근통제, 망간 자료전송 승인·검토 절차, 사용자·단말·위치 기반 접근통제, 이상행위 탐지
Layer 3: 최상위 위험 체급 대상 고도 보안통제 및 신기술 보안 계층	최상위 위험 체급 기관	고도 위험 및 신기술 기반 공격 표면에 대해 강화된 통제, 검증 절차, 개선 조치가 운영되고 있는지 확인	AI 모델 무결성 검증, 프론트 인젝션 방어, OT/ICS망 가시성 확보, 마이크로 세그멘테이션 기반 내부 확산 통제, 레드팀 테스트 또는 공격 시뮬레이션 기반 검증

2. 위험 체급별 평가지표 차등 할당 모델

본 절에서는 4장에서 결정된 3개의 위험 체급(Tier 1, Tier 2, Tier 3)과 1절에서 정의한 3개의 평가 지표 계층(Layer 1, Layer 2, Layer 3)을 상호 매핑하는 논리적 구조를 제시한다. 기존의 실태평가가 모든 기관에 획일적인 지표를 공통 적용함으로써 기관별 위험 차이를 충분히 반영하지 못했다면, 본 할당 모델은 기관의 위험 수준에 따라 평가지표의 범위와 강도를 차등화하는 위험 비례성(Risk Proportionality) 원칙을 따른다(National Institute of Standards and Technology 2012). 또한, 상위 체급의 기관일수록 하위 체급의 보안 요구사항을 기본으로 포함해야 하는 누적 할당 방식을 채택하여 계층형 방어 구조가 단계적으로 적

용되도록 한다. 구체적인 체급별 평가 지표 할당 모델은 다음과 같다.

1) Tier 3(일반 관리군): 기초 위생 중심의 단일 계층 할당 (Layer 1)

- 대상 및 목적: 상대적으로 위험 노출도가 낮고, 외부 연결성과 데이터 민감도가 제한적인 표준 업무 중심 기관 그룹이다.
- 할당 전략: 불필요한 기술적·행정적 부담을 최소화하기 위해 Layer 1(사이버 위생 및 기초 통제) 지표만을 단독 할당한다. 최신 보안 기술의 도입보다는, 자산 관리나 접근 통제와 같은 사이버 보안의 기본 기반에 해당하는 필수 위생 항목들을 충실히 이행하는지 평가의 초점을 맞춘다.

2) Tier 2(중점 관리군): 동적 경계 방어가 추가된 이중 계층 할당

(Layer 1 + Layer 2)

- 대상 및 목적: 대국민 클라우드 서비스를 운영하거나 대량의 민감 정보(S등급)를 취급하는 등, 내부망과 외부망의 경계 변화가 활발하여 중간 이상 수준의 위험에 노출된 기관 그룹이다.
- 할당 전략: Tier 3가 수행하는 기초 위생(Layer 1)을 전제로, 능동적 방어를 위한 Layer 2(동적 경계 및 심화 통제) 지표가 추가로 할당된다. 클라우드 설정(CSPM) 점검, 신원 기반의 제로 트러스트 인증(Rose 외 2020), 망간 전송 구간의 무해화(CDR) 등 확장된 연결성과 데이터 흐름을 통제할 수 있는 역량을 중점적으로 평가한다.

3) Tier 1(최상위 위험군): 완전 계층 할당 (Layer 1 + Layer 2 + Layer 3)

- 대상 및 목적: 하드코딩 트리거를 통해 사전 편입된 12개 핵심 위험 벡터(국가 안보급 DGI C등급 데이터 보유)와 알고리즘 서열화 최

상위 그룹이 병합된 국가 최우선 방어 대상 기관들이다. 이들은 국가 배후 해킹 조직(Nation-state Actors)의 1차 타깃이 된다.

- 할당 전략: 이 그룹에게는 하위 두 계층(Layer 1, Layer 2)은 물론, 최상위 난이도의 Layer 3(고도 보안통제 및 신기술 보안) 지표까지 추가로 할당된다. 최고 수준의 위협을 가진 기관에게 그에 상응하는 강화된 보안통제와 검증 절차를 적용하기 위한 것으로, 레드팀 테스트, 공격 시뮬레이션 기반 검증, AI 모델 무결성 검증, OT/ICS 등 특수망 가시성 확보와 같은 고도 통제 항목을 포함할 수 있다.

3. 비보상적 평가 및 유연한 강등

기존 사이버 보안 실태평가의 주요 한계점 중 하나는 평가 항목 간의 점수 합산 및 만회가 허용되는 ‘보상적 구조’에 있다. 즉, 기관이 핵심 보안 통제 영역에서 중대한 미흡 사항을 보이더라도, 상대적으로 이행이 용이한 타 일반 지표에서 높은 점수를 획득하여 종합 평가 등급을 상향시킬 수 있는 점수 상쇄 현상이 발생할 우려가 있었다. 이는 기관의 실질적인 위협 노출도가 종합 결과에 충분히 반영되지 못하고 핵심 위험이 희석되는 평가 체계의 구조적 한계를 야기한다. 본 연구는 이러한 평균의 함정을 완화하고 평가의 실효성을 높이기 위해, 핵심 지표 간 점수 만회가 제한되는 비보상적 평가 규칙과 이에 기반한 유연한 강등 로직을 제안한다.

1) 계층 간 하향 종속성 기반의 과락 모델

본 모델의 평가 지표는 적층형 방어 구조를 전제로 하므로, 상위 계

층의 보안통제는 하위 계층의 기본 통제가 일정 수준 이상 충족될 때 의미를 가진다. 이를 반영하기 위해 본 연구는 계층 간 하향 종속성(Downward Dependency)에 기반한 과락 모델을 도입한다. 예를 들어, Tier 1 기관이 Layer 3의 고도 보안통제 지표에서 우수한 평가를 받더라도, 그 전제 조건이 되는 Layer 1의 기초 위생 지표에서 중대한 미흡이 확인될 경우, Layer 2 및 Layer 3의 평가 결과를 그대로 인정하지 않고 상위 계층 평가 결과의 반영을 제한하거나 재검토하도록 설계할 수 있다. 이는 기초 보안통제가 충족되지 않은 상태에서 고도 보안통제의 성과만으로 전체 평가 결과를 보완하는 문제를 방지하기 위한 것이다.

2) 핵심 위험 지표 미충족 시 평가 결과 강등

각 위험 체급별로 기관의 핵심 위험을 직접적으로 반영하는 일부 지표를 필수 핵심 지표로 지정할 수 있다. 예를 들어 Tier 1에서는 AI 모델 무결성 검증, 프롬프트 인젝션 방어, 특수망 가시성 확보 등이, Tier 2에서는 망간 자료전송 무해화(CDR) 적용률이나 외부 연계 접근 통제 수준 등이 필수 핵심 지표로 설정될 수 있다. 해당 지표에서 중대한 미흡이 확인될 경우, 다른 지표의 높은 점수와 무관하게 최종 평가 결과를 한단계 아래로 강등(Downgrade)하거나 조건부 판정을 부여하는 방식이 가능하다.

- 평가 결과 강등의 정책적 효과: 단순히 일부 점수를 감점하는 방식보다, 핵심 위험 지표의 미충족이 최종 평가 결과에 직접 반영되도록 함으로써 기관장 및 경영진에게 핵심 보안통제의 중요성을 명확히 전달할 수 있다.
- 유연성의 확보: 이러한 강등 조치는 영구적 불이익이 아니라, 핵

심 미흡 사항의 신속한 보완을 유도하기 위한 절차적 장치로 설계된다. 기관이 해당 필수 핵심 지표의 미흡 사항을 보완하고 재검증을 통과하면, 제한된 평가 결과를 회복할 수 있는 경로를 제공할 수 있다. 이는 평가가 단순한 제재가 아니라 핵심 위험의 신속한 조치와 실질적 보안 수준 향상을 유도하기 위한 장치임을 의미한다.

VI. 결론 및 향후 연구 과제

1. 연구 결과 요약

본 연구는 공공기관 사이버보안 실태평가 체계가 획일적 평가지표와 점수 합산 방식으로 인해 기관별 실질 위험을 충분히 반영하지 못할 수 있다는 문제의식에서 출발하였다. 이에 본 연구는 기관의 위험 수준에 따라 평가지표의 범위와 강도를 차등화하는 위험 체급 기반 사이버보안 실태평가 프레임워크를 제안하였다. 이를 위해 첫째, 데이터 중력(DGI), 망 복잡도(NCI), 방어 역량(ACI)으로 구성된 위험 DNA를 정의하고, 전문가 쌍대 비교와 옐로 레이팅을 활용하여 위험 벡터 간 상대 위험 서열을 산출하였다. 둘째, 산출된 위험 체급에 따라 기초 보안통제, 심화 통제, 고도 보안통제로 구성된 계층형 평가지표를 차등 할당하는 모델을 제시하였다. 셋째, 핵심 보안통제의 미흡이 다른 평가항목의 높은 점수로 상쇄되지 않도록 임계치 과락과 핵심 지표 미충족 시 최종 평가 결과를 강등하는 비보상적 평가 로직을

제안하였다. 이러한 프레임워크는 기관별 위험 특성에 따라 평가 강도와 평가지표 범위를 사전에 조정하고, 핵심 위험에 보안 역량을 집중하도록 유도하는 위험 기반 실태평가 체계의 절차적 기반을 제공한다.

2. 연구의 의의와 정책적 함의

본 연구가 제시하는 프레임워크는 다음과 같은 의의와 정책적 함의를 가진다. 첫째, 본 연구는 위험 체급을 단순한 분류 결과가 아니라 평가 강도, 평가지표 범위, 보안통제 요구 수준을 결정하는 위험 기반 의사결정 기준으로 활용할 수 있다는 점에서 의의가 있다. 이를 통해 상위 위험 체급 기관에 강화된 평가를 적용하는 이유를 명확히 제시하고, 피평가 기관의 정책적 수용성과 승복력을 높일 수 있다. 둘째, 국가 보안 자원의 전략적·효율적 배분 가능성을 제시한다. 모든 기관에 동일한 수준의 평가 부담과 보안 요구사항을 부과하는 대신, 고위험군에는 정밀한 방어 의무와 고도 보안통제를 적용하고, 저위험군에는 기본 보안통제 중심의 평가를 적용함으로써 한정된 보안 예산과 전문 인력을 핵심 위험 방어에 집중시킬 수 있다. 이는 기관의 위험 수준에 상응하는 보안 요구를 부과하는 위험 비례성 원칙의 정책적 구현 가능성을 보여준다. 셋째, ‘평균의 함정’을 완화함으로써 보안 평가의 실효성을 높일 수 있다. 비보상적 평가 로직은 기초 보안 통제 미흡이나 필수 핵심 지표의 실패가 다른 평가항목의 높은 점수로 상쇄되는 문제를 제한한다. 이는 기관이 단순한 종합점수 관리가 아니라, 기초가 견고한 실질적 방어 성숙도를 갖추도록 유도한다는 점에서 의미가 있다. 넷째, 기관의 자율적 공격 표면 축소를 유도

하는 거버넌스 기제로 활용될 수 있다. 기관이 편의를 위해 신기술을 무분별하게 도입하거나 관리되지 않는 자산, 즉 새도우 IT를 확장할 경우, 위험 체급 상승에 따라 더 높은 수준의 평가지표와 보안통제 요구가 부과된다. 결과적으로 각 기관은 기술 도입의 편익과 그에 따른 보안 비용을 함께 고려하게 되며, 이는 기관이 주도적으로 불필요한 외부 연계와 공격 표면을 관리하도록 유도하는 자율적 보안 거버넌스 정착에 기여할 수 있다.

3. 연구의 한계점 및 향후 과제

본 연구는 공공기관 사이버보안 실태평가 2.0 개발 요구사항에 대응하여 위험 DNA 기반 체급화 프레임워크를 제안하는 데 초점을 두었으며, 실제 전문가 패널의 쌍대 비교 데이터를 활용한 경험적 검증까지 수행하지는 못했다. 본 연구에서 활용한 Elo 레이팅과 Fisher-Jenks 알고리즘은 각각 쌍대 비교 결과의 상대 서열화와 1차원 점수 분포의 자연적 단절 분류를 위한 기존 알고리즘이므로, 본 연구의 핵심 한계는 알고리즘 자체의 수학적 작동 여부가 아니라 위험 벡터 간 상대 위험 판단 데이터를 실제 전문가 집단으로부터 수집하지 못했다는 점에 있다.

따라서 향후 연구에서는 현장 보안관리자, 정보보호 책임자, ISMS-P 선임심사원, 공공부문 보안감사 경험자 등 통합적 보안평가 역량을 갖춘 전문가 패널을 구성하여, 36개 비-C등급 위험 벡터에 대한 실제 쌍대 비교 데이터를 수집할 필요가 있다. 이를 통해 DG-I·NCI·ACI 조합에 대한 실제 전문가 집단의 상대 위험 판단 데이터를 확보하고, 이를 본 연구의 위험 체급화 절차에 적용할 필요가 있

다. 또한 전문가 판단의 일관성을 높이기 위해 사전 판단 기준 공유, 델파이 방식의 반복 검토, 평가자 간 일치도 검토 등의 절차를 후속 연구에서 구체화해야 한다.

아울러 Fisher-Jenks 기반 체급 분류 결과를 실제 정책 환경에 적용할 때에는 경계 사례에 대한 검토 기준을 마련할 필요가 있다. Fisher-Jenks 알고리즘은 주어진 위험 서열 점수 분포에서 자연적 단절점을 찾는 방법이므로, 체급 간 구분성을 높이는 데 활용될 수 있다. 다만 실제 제도화 단계에서 체급 경계에 근접한 위험 벡터가 확인되는 경우에는 추가 전문가 검토 또는 보완적 판단 절차를 둘 수 있다.

또한 DGI·NCI·ACI 상태값의 일관된 판정을 지원하기 위한 평가자 가이드 개발이 필요하다. 본 연구는 위험 DNA 메트릭의 기본 정의와 조작적 판단 기준을 제시하였으나, 실제 실태평가 적용 단계에서는 각 메트릭 값을 구체적 사례에 적용하기 위한 세부 설명과 예시가 요구될 수 있다. 향후 연구에서는 CVSS User Guide와 유사하게 각 메트릭 값의 판정 기준, 대표 사례, 경계 사례 발생 시 참고 기준, 필요 증빙자료, 필요시 평가자 간 판단 일관성 검토 절차를 포함하는 위험 DNA 메트릭 평가 가이드를 개발할 필요가 있다.

나아가 본 모델에서 제시한 계층별 평가지표를 실제 평가·감사 절차와 연계하는 방안도 후속 연구 과제로 남는다. 특히 개별 지표의 판정 방식, 감사 절차, 증빙자료 요건은 향후 제도화 단계에서 구체화될 필요가 있다. 또한 자동 수집이 가능한 평가지표는 API 기반 점검·검증 방식과의 연계 가능성을 검토할 필요가 있다.

〈참고문헌〉

- 국가정보원. 2025. 『국가 망 보안체계(N2SF) 보안가이드라인 1.0』.
- 최명길. 2025. “기관 특성을 반영한 사이버보안 실태평가 체계 연구.” 『사이버안보연구』 2집 1호: 45-89.
- 한국인터넷진흥원(KISA). 2026. 『국가 망 보안체계(N2SF) 실증 사례집』.
- 한국사이버안보학회(KACS). 2026. “2026년 실태평가연구회 위탁 연구용역 연구제안 요청서.”
- Center for Internet Security (CIS). 2021. “CIS Community Defense Model 2.0.”
- Elo, A. E. 1978. The Rating of Chessplayers, Past and Present. Arco Pub.
- FIRST. 2023a. “Common Vulnerability Scoring System Version 4.0: User Guide.” Forum of Incident Response and Security Teams.
- FIRST. 2023b. “Common Vulnerability Scoring System Version 4.0: Specification Document.” Forum of Incident Response and Security Teams.
- Jenks, G. F. 1967. “The Data Model Concept in Statistical Mapping.” International Yearbook of Cartography 7, 186-190.
- Louviere, J. J., D. A. Hensher, and J. D. Swait. 2000. Stated Choice Methods: Analysis and Applications. Cambridge University Press.
- National Institute of Standards and Technology (NIST). 2012. “Guide for Conducting Risk Assessments.” NIST SP 800-30 Rev. 1.
- National Institute of Standards and Technology (NIST). 2024. “The NIST Cybersecurity Framework (CSF) 2.0.” NIST CSWP 29.
- Rose, S., O. Borchert, S. Mitchell, and S. Connelly. 2020. “Zero Trust Architecture.” NIST Special Publication 800-207.
- Saaty, T. L. 1980. The Analytic Hierarchy Process. McGraw-Hill.

Dynamic Cybersecurity Posture Assessment Framework based on Risk Tiers

Kim Hong Geun | Graduate School of Information Security, Dongguk University

This study proposes a risk-tier-based dynamic cybersecurity assessment framework for public institutions in response to the relaxation of network separation and the introduction of the National Network Security Framework (N2SF) based on multi-level security. Conventional checklist-based assessments aggregate item-level scores to produce an overall result, which may allow critical security weaknesses to be offset by high scores in other administrative or managerial controls. Such a compensatory structure can dilute core risks and make it difficult to reflect differences in data criticality, external connectivity, and institutional defensive capacity. To address this limitation, this study models institutional cyber risk as a Risk DNA vector composed of Data Gravity Index (DGI), Network Complexity Index (NCI), and Administrative Capacity Index (ACI). Critical risk factors, such as the possession of confidential data, are first assigned to the highest risk tier. The remaining risk vectors are converted into one-dimensional relative risk rankings through expert pairwise comparisons and the Elo rating system. The Fisher-Jenks algorithm is then applied to classify the vectors into three risk tiers, and layered assessment indicators are differentially assigned according to the resulting tiers. In addition, non-compensatory evaluation logic, including minimum threshold failure and critical indicator failure, is applied to reduce the dilution of core risk factors. The proposed framework contributes to a risk-based assessment approach by enabling the scope and intensity of cybersecurity evaluation to be adjusted in advance according to institutional risk characteristics.

Keyword N2SF, Cybersecurity Posture Assessment, Risk DNA, Elo Rating Algorithm, Risk Tier