



DOI 10.28925/2663-4023.2026.34.1356

УДК 004.72; 004.7; 004.8

Сиротинський Роман Михайлович

асистент кафедри захисту інформації

місце роботи: Національний університет «Львівська політехніка», Львів, Україна

ORCID: 0009-0002-6280-3290

roman.m.syrotynskyi@lpnu.ua

Тишик Іван Ярославович

кандидат технічних наук, доцент кафедри захисту інформації

місце роботи: Національний університет «Львівська політехніка», Львів, Україна

ORCID: 0000-0003-1465-5342

ivan.y.tyshyk@lpnu.ua

**ПІДВИЩЕННЯ ВІДПОВІДНОСТІ ПРИНЦИПАМ НУЛЬОВОЇ ДОВІРИ В
АРХІТЕКТУРІ КОРПОРАТИВНИХ МЕРЕЖ ІЗ ВИКОРИСТАННЯМ ШТУЧНОГО
ІНТЕЛЕКТУ**

Анотація. Сучасні корпоративні мережеві інфраструктури характеризуються постійним зростанням кількості інформаційних сервісів, активним використанням гібридних і хмарних середовищ та суттєвим ускладненням політик мережевого доступу. Динамічність таких середовищ призводить до регулярних змін конфігурацій засобів мережевого захисту, тоді як вимоги до взаємодії корпоративних сервісів можуть зберігатися та актуалізуватися в різних інформаційних системах незалежно від фактичної конфігурації мережі. У результаті з часом виникають розбіжності між задекларованими вимогами та реально наданими мережевими доступами, що ускладнює підтримання належного рівня захищеності та відповідності принципам архітектури нульової довіри. Особливо актуальною ця проблема є для великих корпоративних середовищ, де кількість правил мережевого доступу може сягати кількох тисяч, а їх періодичний аналіз та ресертифікація є тривалими та потребують значних часових затрат і людських ресурсів. Традиційні засоби аналізу конфігурацій дозволяють виявляти окремі аномалії, однак не завжди можуть визначити наскільки фактично надані права доступу відповідають їхньому початковому призначенню та актуальним вимогам безпеки корпоративних сервісів. Це створює передумови для накопичення застарілих, надлишкових, незадокументованих або надмірно широких прав доступу та ускладнює реалізацію принципів найменших привілеїв, явної авторизації та постійної перевірки безпекових правил. У роботі досліджується можливість використання технологій штучного інтелекту для підвищення ефективності процесів контролю та ресертифікації політик мережевого доступу в контексті реалізації принципів нульової довіри. Запропоновано концептуальний підхід, орієнтований на автоматизоване опрацювання різномірної інформації про корпоративні сервіси та фактичний стан політик мережевого захисту з метою оцінювання їх відповідності вимогам безпеки та принципам нульової довіри. Описаний підхід передбачає використання результатів інтелектуального аналізу конфігурацій для виявлення потенційних відхилень, ідентифікації правил мережевого доступу, що потребують додаткової перевірки, та формування рекомендацій для відповідальних фахівців. При цьому штучний інтелект розглядається як допоміжний механізм підтримки прийняття рішень, а не як автономний засіб внесення змін до політик безпеки. Застосування такого підходу дозволить скоротити обсяг ручної роботи під час ресертифікації, підвищити масштабованість процесу контролю та сприятиме підтриманню відповідності мережесхем політик принципам нульової довіри у складних корпоративних інфраструктурах.

Ключові слова: політики безпеки, ресертифікація правил, мережевий файрвол, нульова довіра, принцип найменших привілеїв, штучний інтелект.

ВСТУП

Еволюція корпоративних інфраструктур та цифрова трансформація інформаційних систем суттєво видозмінила спосіб проектування, розгортання та експлуатації корпоративних мережесхем інфраструктур.



Кількість корпоративних сервісів зростає а їх архітектура все частіше стає розподіленою. Застосовуються як розділені територіально локальні датацентри так і хмарні, мультихмарні та гібридні платформи. Локальна інфраструктура сегментується та поєднується з публічними хмарними платформами та рішеннями «програмне забезпечення як послуга» (SaaS) для покращення масштабованості, гнучкості та безперервності бізнесу. Одночасно, віддалені та гібридні моделі роботи стали стандартною операційною практикою, що вимагає безпечного доступу до корпоративних ресурсів практично з будь-якого місця та пристрою.

Основний виклик кібербезпеці поступово змістився від захисту чітко визначеного периметра мережі до управління дедалі складнішою екосистемою ідентифікаційних даних, додатків, послуг та політик доступу до мережі. Замість того, щоб лише розширювати поверхню атаки, організації повинні справлятися з операційним навантаженням, пов'язаним з підтримкою тисяч правил брандмауера, політик контролю доступу, конфігурацій VPN, груп безпеки хмари та винятків, специфічних для програм, протягом усього життєвого циклу інфраструктури. Нещодавні дослідження підкреслюють, що суміщення гібридних хмарних середовищ та Zero Trust значно збільшує операційну складність, роблячи ручне управління політиками безпеки дедалі менш практичним [1].

Згідно NIST - Zero Trust передбачає, що активам або обліковим записам користувачів не надається неявна довіра виключно на основі їхнього фізичного або мережевого розташування, як це було прийнято в периметральних моделях. Архітектура Zero Trust (ZTA) забезпечує постійну перевірку кожного запиту на доступ та застосовує принцип найменших привілеїв незалежно від мережевого розташування. Тим не менш, успішне впровадження Zero Trust вимагає постійної оцінки величезної кількості операційних даних, що надходять від брандмауерів, постачальників ідентифікаційних даних, систем захисту кінцевих точок, хмарних платформ, рішень SIEM та інструментів моніторингу мережі. У великих корпоративних середовищах ручне забезпечення узгодженості між задокументованими бізнес-вимогами та фактичними політиками доступу до мережі стає дедалі складнішим, оскільки інфраструктури розвиваються, а нові послуги постійно впроваджуються.

Штучний інтелект (ШІ), зокрема моделі великих мов (LLM) та агентні рішення на основі ШІ пропонує перспективні можливості щодо вирішення цієї проблеми. Здатність аналізувати документи та мережеві конфігурації дозволяють використовувати великі обчислювальні можливості штучного інтелекту на нестандартизованих даних. ШІ має потенціал значно підвищити ефективність перевірки відповідності Zero Trust. Замість того, щоб замінювати адміністраторів безпеки, ШІ служить інтелектуальним механізмом підтримки прийняття рішень щонайменше.

Аналіз останніх досліджень і публікацій. Перехід від традиційної периметрової моделі захисту мережі до архітектури нульової довіри (Zero Trust Architecture, ZTA) розглядається як відповідь на зростання складності та розподіленості сучасних корпоративних інфраструктур. Нульова довіра передбачає відмову від неявної довіри на користь постійної перевірки користувачів, пристроїв та запитів на доступ. Безперервна перевірка є одним із центральних елементів ZTA, а інтеграція штучного інтелекту (AI) та машинного навчання (ML) є перспективним механізмом адаптивного захисту та виявлення загроз [2].

Практична реалізація принципів Zero Trust на мережевому рівні значною мірою залежить від коректності та актуальності політик мережевої безпеки. Автор іншого дослідження визначає «зростаючу складність управління політиками безпеки, особливо в корпоративних мережах» як суттєву проблему та зазначає, що політики часто конфігуруються вручну й ізольовано різними адміністративними групами, що створює передумови для виникнення невідповідностей [3].

Особливо актуально ця проблема є для корпоративних брандмауерів, конфігурації яких можуть містити тисячі взаємопов'язаних правил. Управління великими наборами правил це «складний і схильний до помилок» процес що потребує формального механізму їх аналізу [4]. В дослідженнях відзначається складність управління конфігураціями брандмауерів і пропонується автоматизований підхід до виявлення та усунення аномалій у правилах [5].

Окремим напрямом є використання AI для підвищення адаптивності Zero Trust. В дослідженні авторів з Нігерії для архітектури нульової довіри пропонується AI-enabled Zero Trust framework, у якому AI та ML застосовуються для «адаптивної автентифікації, виявлення аномалій у реальному часі та динамічного контролю доступу», демонструючи потенціал інтелектуальної автоматизації для підтримки принципу найменших привілеїв [6].

Таким чином, наявні дослідження охоплюють Zero Trust, автоматизований аналіз правил брандмауерів та використання AI у кібербезпеці, однак ці напрями переважно розглядаються окремо. Значно менше уваги приділено семантичній перевірці відповідності фактичних правил мережевого доступу задокументованим вимогам корпоративних сервісів. Саме ця прогалина обґрунтовує запропонований у роботі підхід до використання AI для співставлення сервісної документації з

фактичними політиками брандмауерів та підтримки їх періодичної ресертифікації відповідно до принципів Zero Trust.

Особливості архітектури корпоративних мереж сьогодні. Типова архітектура корпоративної мережі є мікросегментованою на численні ізольовані підмережі та складається з кількох взаємопов'язаних шарів. Корпоративні сервіси, їх компоненти та сегменти з іншими типами призначення логічно розділяються на мережевому рівні, доступ між якими контролюється засобами корпоративних файрволів. Користувачі автентифікуються через платформи ідентифікації, такі як Active Directory або Microsoft Entra ID. Доступ може проходити через рішення VPN або Zero Trust Network Access, кластери брандмауерів, демілітаризовані зони, внутрішні сегменти мережі, хмарні середовища, платформи SaaS та архітектури додатків на основі мікросервісів. Спрощений шлях зв'язку можна представити ланцюжком, зображеним на рисунку 1.

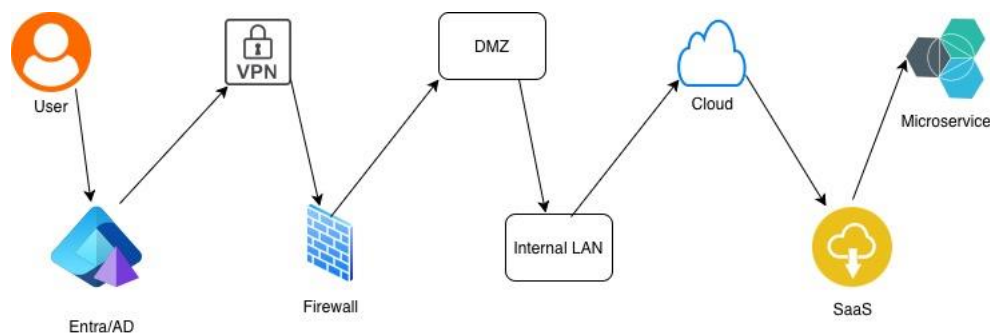


Рис. 1. : Шлях підключення користувача до системи в сучасній інформаційній інфраструктурі

Сучасна мережева архітектура є громіздкою та динамічною. Розгортаються нові сервіси, додатки мігрують на хмарні платформи, додаються інтеграції SaaS, бізнес-підрозділи запитують тимчасовий доступ, а гібридне підключення розширюється в кількох регіонах та постачальниках. Зміна безпекового стану вузлів в мережі передбачає постійний його контроль та потребу адаптивного підлаштування політик доступу під зміну контексту мережевого підключення.

Зростає складність операційної підтримки та супроводу мережево-безпекової інфраструктури корпоративних сервісів, оскільки велика кількість точок застосування політик та контроль актуальності правил з бізнес задачами які ставить власник сервісу можуть йти врозріз один одному. Інформація про безпеку підприємства та його компонентів розподілена між багатьма джерелами та супроводжується різними відповідальними особами. Описи послуг можуть знаходитися на сторінках SharePoint або Wiki. Метадані власності та критичності для бізнесу можуть існувати в CMDB. Запити на зміни можуть зберігатися в ServiceNow. Інформація про ідентифікацію та групи може керуватися в AD або Entra ID. Хмарні активи можуть бути описані в Azure або інших хмарних інвентаризаціях. Однак конфігурації брандмауера підтримуються на окремих платформах безпеки та використовують політики безпеки які безпосередньо не відповідають документації бізнес-послуг або не встигають коригуватися в процесі зміни їхнього актуальності під час свого життєвого циклу.

Це розділення створює розрив в управлінні. Документація описує, який зв'язок має бути дозволений, тоді як правила мережевого файрвола визначають, який зв'язок фактично дозволений. Ці два представлення часто розвиваються окремо. Попередні дослідження управління мережевими політиками показали, що конфігурації контролю доступу можуть містити невідповідності, а управління розподіленими політиками безпеки мережі є складним та схильним до помилок [7]. У подібних роботах з абстракції політик стверджується, що перетворення високорівневих організаційних політик на конфігурації рівня пристроїв є виснажливим, трудомістким та вразливим до людських помилок [8].

Як наслідок, підприємствам потрібні автоматизовані механізми, які можуть пов'язувати наміри обслуговування на бізнес-рівні з технічним конфігураціями засобів мережевого контролю доступу та іншими безпековими конфігураціями.

Прогалини в безпеці як наслідок інфраструктурної складності. Різноманіття платформ та гібридні інфраструктури розширюють потребу в експертизі щоб ефективно їх обслуговувати, консистентність правил доступу потребує розуміння всіх платформ які є в ланцюжку доступу. Мікросегментована інфраструктура збільшує кількість точок застосування політики в рази, а потреба супроводу доступів протягом їх життєвого циклу в силу динамічності внутрішніх і зовнішніх середовищ



збільшує кількість необхідних ресурсів необхідних для обслуговування великої кількості політик безпеки та виконувати поставлені на мережеву інфраструктуру бізнес завдання.

Зі зростанням корпоративних мереж бази правил файрволів, як правило, накопичують історичні та операційні артефакти. Типовими прикладами є застарілі правила для виведених з експлуатації сервісів, тимчасові правила, які ніколи не видалялися бо про них забули, широкі правила «Дозволити будь-які», дублікати або тіньові правила, недокументовані шляхи доступу, правила без відомого власника та правила які не використовуються (по яких не проходить трафік). Ці проблеми є не лише операційною неефективністю; вони безпосередньо впливають на рівень захищеності інфраструктури та послаблюють впровадження принципів нульової довіри.

Великі масиви правил в конфігураціях можуть спричиняти аномалії доступу. Дослідження нагромадження політик брандмауера виявляють такі проблеми, як надмірність, суперечності, нерелевантні правила та правила, що посилаються на неіснуючі хости або служби.[9] Дослідження конфігурацій безпеки розподілених мереж також показують, що невідповідності між брандмауерами, системами виявлення вторгнень та іншими компонентами безпеки можуть погіршити щільність безпекових заходів.[10]

Правила міжмережевих екранів які відкривають широкий доступ є особливо проблематичними. Відкриття занадто широких діапазонів source IP, непотрібні порти чи аплікації призначення або недокументовані адміністративні протоколи, можуть надати зловмисникам додаткові шляхи для горизонтального руху. Це суперечить припущенню Zero Trust (нульової довіри), що жодному внутрішньому сегменту мережі не слід довіряти за замовчуванням. Дослідження по використанню систем візуалізації політик брандмауера підкреслюють, що надмірно дозвольні політики та некеровані політики можуть призвести до ненавмисного виставлення внутрішніх ресурсів та систем в інтернет. [11] У попередніх роботах з ідеями візуалізації правил брандмауера також зазначається, що складність політик ускладнює безпечну оптимізацію та збільшує ризик хибних або небезпечних налаштувань правил. [12]

З точки зору Zero Trust, ці прогалини порушують три основні принципи. По-перше, вони порушують принцип найменших привілеїв, оскільки доступ може бути ширшим, ніж вимагає сервіс. По-друге, вони послаблюють безперервну перевірку, оскільки правила можуть залишатися активними ще довго після того, як початкова бізнес-потреба зникла. Або банальна вразливість до спуфінгу чи підставляння дозволених айпі адрес. По-третє, вони ставлять під загрозу явну авторизацію, оскільки недокументовані правила брандмауера не можуть бути чітко пов'язані із затвердженими вимогами до сервісу або записами про зміни. Це створює дрейф політик: поступову розбіжність між задокументованим наміром або його відсутністю та реальною конфігурацією мережі.

Ручний перегляд наборів правил може зменшити ці ризики, але ручний аналіз погано масштабується у великих мікросегментованих та гібридних мережеских архітектурах. Дослідження управління політиками брандмауера неодноразово показували, що виявлення та вирішення аномалій є складними, коли політики містять тисячі правил і задіяні численні адміністратори. [13]

Процес ресертифікації мережевого доступу є елементом життєвого циклу правил доступу та згідно ідеології нульової довіри повинен виконуватися на періодичних засадах, проте на практиці через динамічне зростання об'єктів конфігурації, погано налагоджену процесну компоненту та високу ресурсозатратність таких процедур де кількість правил вимірюється десятками тисяч - така синхронізація часто не відбувається, правила створюються і більше ніколи ніхто їх не переглядає.

Опис проблеми та мета статті. Успішне впровадження архітектури нульової довіри (ZTA) залежить від підтримки постійної узгодженості між задокументованими бізнес-вимогами та фактичними політиками мережевої безпеки, що застосовуються в корпоративних інфраструктурах. У сучасних організаціях вимоги до мережевого доступу зазвичай описуються в документації до служб, базах даних керування конфігураціями (CMDB), архітектурних репозиторіях та системах управління знаннями, таких як SharePoint або ServiceNow. Водночас мережевий зв'язок забезпечується за допомогою наборів правил брандмауера, груп хмарної безпеки та інших механізмів контролю доступу, які постійно розвиваються протягом життєвого циклу інфраструктури.

У міру розширення корпоративної інфраструктури ці два інформаційні представлення можуть еволюціонувати незалежно одне від одного. Документація до сервісів відображає передбачувану модель взаємодії між бізнес-застосунками та інформаційними ресурсами, тоді як конфігурації брандмауерів, хмарних механізмів контролю доступу та інших засобів мережевого захисту змінюються відповідно до розвитку інфраструктури, її міграції до хмарного середовища, реалізації тимчасових запитів на доступ та внесення змін артефактів у процесі реагування на інциденти безпеки. Унаслідок такої незалежної еволюції з часом можуть виникати недокументовані шляхи мережевої взаємодії, застарілі правила брандмауерів, дубльовані або надлишкові політики, а також надмірні права доступу, які вже не відповідають актуальним бізнес-вимогам і фактичним потребам сервісів.



Основна проблема полягає не лише у виявленні потенційно небезпечних чи неактуальних правил брандмауера, а й у забезпеченні постійної перевірки відповідності фактично впроваджених мережових політик задокументованим вимогам до взаємодії корпоративних сервісів та принципам архітектури нульової довіри. У великих корпоративних мережах, які містять тисячі правил доступу та численні незалежні джерела технічної і сервісної документації, процес такої ресертифікації потребує значних часових затрат і людських ресурсів, що унеможливує його регулярне виконання в ручному режимі. У результаті ускладнюється своєчасне виявлення відхилень між задекларованими вимогами та фактичною конфігурацією мережі, що порушує фундаментальні принципи Zero Trust: найменших привілеїв (Least Privilege), безперервної перевірки (Continuous Verification) та явної авторизації (Explicit Authorization). Отже, задача цієї роботи, полягає у підвищенні ефективності та масштабованості процесу ресертифікації мережових правил на предмет їх відповідності корпоративним вимогам і концепції Zero Trust.

Мета статті полягає в описі розробки підходу до автоматизованої періодичної семантичної перевірки відповідності фактично реалізованих політик мережевого доступу задокументованим вимогам та принципам архітектури нульової довіри. Запропонований підхід має забезпечувати аналіз і зіставлення інформації з корпоративної сервісної документації та конфігурацій брандмауерів, виявлення розбіжностей між задекларованими й фактично реалізованими мережевими доступами, а також ідентифікацію надлишкових, застарілих, невикористовуваних або надмірно дозвільних правил. Результатом перевірки має бути оцінка відповідності мережових політик корпоративній документації та концепції Zero Trust, зокрема її ключовим принципам – найменших привілеїв і безперервної перевірки, з формуванням обґрунтованих рекомендацій щодо їх актуалізації або ресертифікації.

Штучний інтелект як засіб забезпечення дотримання принципу нульової довіри. Штучний інтелект може підтримувати дотримання принципів нульової довіри, забезпечуючи безперервний, масштабований та контекстно-залежний аналіз мережових політик. У запропонованому підході Штучний інтелект не приймає остаточних рішень щодо зміни конфігурації. Натомість він допомагає адміністраторам зрозуміти, чи відповідають існуючі правила брандмауера задокументованим вимогам до послуг. Це узгоджується з нещодавнім дослідженням нульової довіри яке позиціонує ШІ як механізм виявлення аномалій, контекстного аналізу, адаптивного контролю доступу та швидшого реагування, а не як повну заміну управління [14].

Важливими внесками ШІ в цьому випадку використання є:

Збільшення продуктивності та відсутність людського фактору при аналізі політик брандмауера на наявність технічних аномалій, таких як дублікати, перекриття правил, конфлікт, тіньові правила, правила без спрацювань.

Семантична інтерпретація. Порівняння та аналіз існуючих правил безпеки в конфігурації міжмережових екранів з номіналом, описаним в документації сервісу як частина ресертифікації в життєвому циклі мережевого доступу.

Моделі великих мов можуть обробляти неструктуровані описи послуг, допоміжні документи, сторінки архітектури, запити на підтримку, та операційні нотатки. Вони можуть витягувати структуровані вимоги до зв'язку, такі як підмережа джерела, система призначення, протокол, порт, власник послуги, середовище, бізнес-мета та термін дії. Дослідження управління доступом на основі штучного інтелекту також підкреслюють важливість аналізу поведінки користувачів, стану пристрою, місцезнаходження та контекстуального ризику для підтримки рішень щодо безперервного доступу [15].

ШІ також може підтримувати класифікацію та пріоритезацію аномалій в правилах міжмережових екранів. У контексті дотримання Zero Trust цей аудит може ґрунтуватися не лише на технічному перекритті, але й на сигналах ризику, таких як недокументований доступ, адміністративні порти, необмежений доступ з або в інтернет, широкі діапазони джерел, відсутні метадані власника або невідповідність із затвердженою документацією служби.

ШІ може додатково виявляти прогалини між передбачуваним та фактичним доступом. Наприклад, якщо в документації зазначено, що служба ERP вимагає доступу HTTPS з підмережі Accounting до певного сервера ERP, але правило брандмауера дозволяє будь-якому джерелу отримувати доступ до SSH, HTTPS та RDP на цьому сервері, ШІ може класифікувати правило як надмірно дозвільне та частково недокументоване. Це підтримує безперервну перевірку шляхом порівняння реального стану політики із задокументованими бізнес-намірами. Поведінкові та довірчі підходи до управління доступом за принципом Zero Trust також показують, що статистичні моделі та моделі машинного навчання можуть допомогти автоматизувати елементи підтримки автентифікації та авторизації в умовах Zero Trust [16].

Однак, штучний інтелект має залишатися механізмом підтримки рішень. Агенти можуть галюцинувати, неправильно розуміти застарілу документацію або неправильно робити висновки про технічні вимоги. І при тому всьому аргументувати свої рішення, що може виглядати реалістично та впевнено. Тому запропонована модель має надавати пояснення своїм рекомендаціям, посилання та докази

з вихідною документацією та очікувати схвалення адміністратора перед впровадження будь-яких змін в живій конфігурації.

Концептуальна модель семантичної перевірки відповідності правил мережевого доступу. Ідея семантичної перевірки відповідності полягає в створенні моделі яка б реалізовувала рішення проблематики описаної в постановці проблеми, а саме перевірка фактичної конфігурації мережевих засобів захисту інформації на предмет відповідності з наступними наміреннями:

- Документація - вимоги інформаційних сервісів які послугуються корпоративними мережами як інфраструктурним сервісом
- Відповідність принципам нульової довіри
- Актуальний стан доступу згідно його життєвий цикл

В кожному підприємстві по факту існує дві мережеві моделі. Одна мережева модель - намірення. Це те які очікування є у компанії щодо безпекової моделі її інфраструктури. Вона відображає наступне:

- Найкращі безпекові практики
- Висока доступність та продуктивність
- Помірна ресурсозатратність в побудові та обслуговуванні
- Рівень дозволів відповідає мінімальному доступу необхідному для функціонування корпоративних сервісів

Дана модель описана в корпоративних документах, сторінках wiki, в системах управління конфігураціями, сервісних каталогах архітектурній документації та в засобах автоматизації бізнес процесів. Другою є впроваджена мережева модель. Це та конфігурація мережевих пристроїв і правила доступу які є фактично присутніми та впровадженими. Реальність, на противагу очікуванням. Це те як мережа по факту працює. Наприклад набір політик доступу до інтернету на корпоративному міжмережевому екрані, елемент якого може виглядати наступним чином:

```
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS from DMZ
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS to Internet
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS source any
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS destination any
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS application ssl
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS service
application-default
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS action allow
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS log-end yes
set device-group CORP-INET pre-rulebase security rules DMZ-to-Internet-HTTPS profile-setting
group Outbound-Security-Profile
```

З часом документація може розходитися з реальною конфігурацією і саме тут і виникає проблема, коли ніхто не перевіряє чи запланована модель відповідає реалізованій. Для її вирішення пропонується концептуальну модель семантичної перевірки відповідності що складається з 5 блоків, а саме:

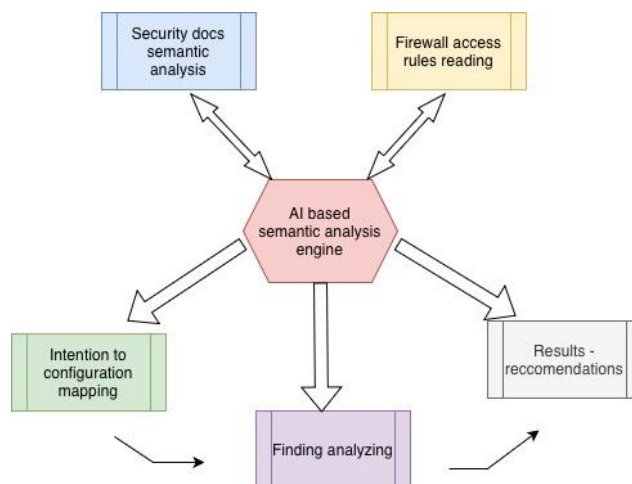


Рис. 2. Структура концептуальної моделі перевірки відповідності

Блок 1. Семантичний аналіз запланованих доступів. Безпекова модель підприємства визначає рівень доступності її корпоративних сервісів та їх компонентів в межах корпоративної мережі. Намірення і вимоги фіксуються в корпоративних політиках та сервісній документації які зберігаються окремо від мережевої інфраструктури в виді документів. Такі задокументовані відомості є першоджерелом інформації яка має трансформуватися в політики доступу міжмережевих екранів в місцях їхнього застосування. Ініціюють та регулюють необхідні доступи працівники відділу захисту інформації та власники інформаційних сервісів, системи яких послугуються корпоративною мережею. На основі задокументованих даних створюються запити на конфігурацію щодо впроваджується змін до політик міжмережевих екранів та інших засобів контролю мережевого доступу. Завдання семантичної моделі на даному етапі є читання корпоративних документів та їх аналіз на предмет опису очікуваних від мережі доступів. Агент аналізує доступні документи та витягує структуровані елементи політик доступу, такі як: джерело, призначення, сервіс, порт, умови доступу. Цінність використання AI на даному етапі відображається в автоматизації процесу розпізнавання нестандартизованих даних, нормалізації та структуруванні результатів читання. Використання агентів дозволяє аналізувати документацію з високою продуктивністю, низькою ймовірністю помилки через причину людського фактору та високою толерантністю до неоднорідності стилів представлення інформації.

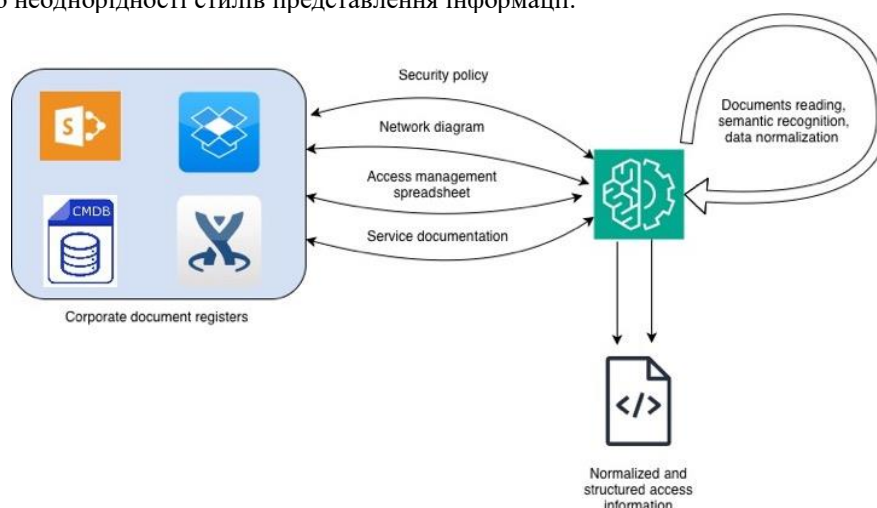


Рис. 3. Структура читання та нормалізації документів мережевої безпеки

Реалізація даного етапу передбачає побудову місточків зв'язності між середовищем LLM та місцем зберігання документації. На практиці, документація корпоративних послуг може бути отримана за кількома сценаріями, починаючи від ручного надсилання документів і закінчуючи автоматизованими конекторами з репозиторіями SharePoint, CMDB, ITSM та Infrastructure-as-Code. Це дозволяє системі штучного інтелекту порівнювати передбачувані повідомлення про доступи, задокументовані в корпоративних репозиторіях, з фактичними правилами брандмауера, зберігаючи при цьому контроль адміністратора над остаточними рішеннями в випадку розходження.

На даному етапі після вчитки документів має відпрацювати ШІ парсер документації. Це компонент, який читає неструктуровані або напівструктуровані документи і витягує з них технічний зміст. Його задача - перетворити людську документацію на структуровані вимоги до мережевого зв'язку. (network communication requirements.)

Функції роботи парсера описані в таблиці 1. Важливо: ШІ не повинен просто “переказувати документ”. Він має формувати нормалізований запис доступу.

Таблиця 1

Функція	Приклад
Entity extraction	знаходить назви сервісів, серверів, subnet, ports
Relationship extraction	визначає, хто з ким має комунікувати
Protocol inference	“HTTPS” → TCP 443
Owner extraction	“Finance IT owns ERP” → owner = Finance IT
Environment detection	production / test / dev
Validity extraction	“temporary until 31.12.2026”
Evidence linking	зберігає посилання на документ і фрагмент тексту
Confidence scoring	оцінює, наскільки чітко описано доступ



Блок 2. Читання та нормалізація правил файрволів. Даний блок семантичної моделі складається з двох основних компонентів, перший з яких це колектор правил доступу. Задача колектора - збирати фактичні правила з firewall-систем і безпекових контролів хмарних оточень. Ціль - визначити, а що реально дозволено на мережі. Джерелами виступають корпоративні файрволи, хмарні файрволи, хмарні безпекові групи, списки доступу, безпекові політики Kubernetes, а також трафік логи, лічильники спрацювань, тощо. Політики доступу вчитуються з різнотипних платформ та продуктів від різних постачальників. Формат та синтаксис в кожному випадку будуть різними. Невід'ємним етапом є пошук та вибір шляхів передачі конфігурації від мережових систем до агента III. Варто розглядати шляхи не як один "конектор до firewall", а як модуль передачі даних для III базованої системи відповідності (ingestion layer for AI-based compliance framework). Поширеними є наступні сценарії:

- Запит конфігурації в режимі "лише читання" з допомогою інтерфейсу API. Це найкращий сценарій в виробництві для класичних корпоративних брандмауерів. AI-агент або проміжний колектор підключається до інтерфейсу управління через API у режимі "лише читання".
- Експорт правил в файл. Це простіший і безпечніший варіант. Адміністратор систем безпеки регулярно експортує набори правил у виді CSV, JSON, XML або в якомусь пропрієтарному форматі, а AI-агент аналізує цей файл.
- Телеметрично базована передача конфігурації файрвола. Сховище правил показує, що дозволено. А SIEM/logs показують, що реально використовується. AI-агент може отримувати не тільки конфігурацію, а й певну мережево-безпекову телеметрію яка доступна в сучасних міжмережових екранах.
- Інтеграція даних через хмарний API. Для хмарних і гібридних інфраструктур важливо читати не тільки класичні файрволи, а й хмарні стандартні системи контролю доступу. До таких відносяться NSG, Azure Firewall, Application Security Groups, Private Endpoints, Route Tables, NACLs, VPC route tables, AWS Network Firewall, Transit Gateway policies та інші.
- Аналітичне зчитування конфігурації з бекапів. У багатьох корпоративних інфраструктурах вже є системи, які регулярно знімають бекапи конфігурацій мережових пристроїв. Це хороший шлях, якщо компанія не хоче давати агенту III прямий API-доступ до інтерфейсів керування міжмережових екранів.
- III доступ на базі МСП інструментів. Актуальний варіант для AI-agent архітектур — дати агенту не прямий доступ до систем, а контрольовані утиліти. Наприклад AI-агент має доступ до команд типу show і read, проте не має доступу до команд write чи remove. Даний сценарій має також потенціал оскільки всі великі виробники мережових продуктів пропонують вже готову МСП сутність для контрольованого підключення до їх продуктів.

В межах другого блоку семантичної моделі для подальшого аналізу та співставлення необхідною є нормалізація правил до єдиного формату та структури. Таким чином вчитані різного синтаксису і платформи правила аналізуються та записуються в стандартизованому форматі для подальшого опрацювання.

З міркувань безпеки AI-агенту краще передавати не "сирий повний доступ до файрволу", а нормалізовану, мінімізовану і контрольовану копію його конфігурації.

Блок 3. Встановлення відповідності значення в документації з фактичним правилом. В центрі семантичної моделі знаходиться блок порівняння намірень, описаних в документації та фактичних налаштувань, застосованих в засобах мережевого контролю. Це якраз ключова частина, де система вже має та розуміє дані з обох сторін і мапить їх між собою. Результатом даного процесу є таблиця з кореляцією даних, отриманих з 2 джерел які доповнюються відповідними статусними позначками, такими як повне чи часткове співпадіння, або відсутність відповідника.

Даний і наступний блоки це "мозок" порівняння. Вони відповідають на питання: "Чи відповідає реально правила файрвола задокументованому наміру?".

Блок 4. Аналізатор даних відповідності. Сирі дані отримані на попередньому етапі вказують на співпадіння але ще не дають оцінку і відповідно і цінність факту присутності чи відсутності відповідника. Наступним кроком є аналіз та висновки, які випливають з співставлення. Даний етап сфокусований більше на аналізі різниці між тим що є задокументовано та фактом, присутнім в налаштуваннях, з додаванням додаткового безпеково мережевого контексту. III базований аналізатор швидко та продуктивно аналізує великі масиви даних, це дозволяє даний процес проводити частіше, уникаючи ресурсних перевитрат в порівнянні з ручним аналізом.

На основі проаналізованих даних, побудована система формує звіт з результатами. Можливими є наступні висновки проведеного аналізу:

- Повне співпадіння. Документація і правило міжмережевого екрану повністю збігаються.



- Часткове співпадіння. Не всі умови доступу співпали, присутні відмінності. Додаються уточнення які саме.
- Правило відсутнє в документації але присутнє на файрволі. Це 100% проблема, оскільки немає підтвердження легітимності даного правила, факт тіньової конфігурації, безпекові ризики.
- Правило присутнє в документації але відсутнє на файрволі. Такий факт це також проблема, що заключається в ризиках обмеженої доступності та може мати вплив на сервіс чи аплікацію для якої створювався запит на доступ.
- Правило доступу дозволяє більше ніж задокументовано. Різновид висновку - пункт 2 з фактом розширення доступу в правилі. Відхід від принципу найменших привілеїв.
- Застаріле правило. Старе тимчасове правило, або правило з протермінованим терміном придатності.
- Правило яке не використовується. Є відповідність в документації, проте протягом періоду який більший за якийсь умовний термін відсутні спрацювання в лічильниках. Питання ресертифікації правила та валідації з власником.
- Конфліктне правило. Ситуація де доступ описується в різних точках документації, де в одному місці написано дозволити, а в іншому - заборонити.

Додаткова цінність результатів аналізу задокументованих та фактичних безпекових політик полягає в потребі відповідати вимогам сучасним безпековим стандартам, в можливості використання їх при періодичному перегляді мережових доступів як елемент їхнього життєвого циклу. Використання штучного інтелекту дозволяє проводити аналітичні сесії з високою інтенсивністю та не витрачати значних людських ресурсів, навіть у випадку великих розподілених інфраструктур з десятками тисяч політик.

Блок 5. Рекомендації. На завершальному етапі система генерує рекомендації для операторів щодо виявлених зауважень. Це хороша стартова точка для прийняття рішення. Ціль не вносити зміни автоматично, а сформулювати пропозицію на перегляд адміністраторами. Рекомендації складаються не лише з дій але й надають пояснення та аргументацію чому саме такий висновок, описують можливий контекст та пропонують декілька варіантів вирішення знайденої проблеми. Разом з результатами аналізу якісна рекомендація має містити:

- Що знайдено ?
- Чому це проблему ?
- Який принцип Zero Trust порушено ?
- Які документи були використані ?
- Яка запропонована дія ?
- Чи потрібна перевірка власника ?
- Який очікуваний ефект ?

В залежності від знайденої проблеми, можуть рекомендуватися сценарії виправлення, що описані в таблиці 2

Таблиця 2

Ідентифікована проблема	Сценарій вирішення
Занадто широке правило	Звузити адреса джерел
Забагато портів чи аплікацій	Видалити порти які не задокументовані
Відсутній власник	Призначити власника доступу
Відсутня документація	Замовити документацію або вимкнути правило
Правило яке не використовується	Запланувати видалення після валідації
Тимчасове правило завершеним терміном придатності	Видалити або перепогодити
Тіньове правило	Видалити або підняти
Відсутнє правило яке мало бути	Створити запит на впровадження якщо бізнес підтвердить необхідність
Дублікати	Об'єднання

РЕЗУЛЬТАТИ ДОСЛІДЖЕННЯ ТА РЕАЛІЗАЦІЯ РІШЕННЯ

Результатом роботи даного дослідження є робочий прототип системи на базі концептуальної моделі перевірки відповідності мережових правил в корпоративній мережі. Рішення здатне опрацьовувати живі дані та формувати результати аналізу правил міжмережових екранів.

Реалізація рішень на базі запропонованої моделі передбачає застосування та інтеграцію ШІ асистентів в інформаційну інфраструктуру та системи корпоративного документообігу. В якості



експерименту використано III асистент Claude Enterprise. Доступ до аналізу документації здійснюється з використанням вбудованого конектора Microsoft365. Даний конектор дозволяє підключитися до корпоративного середовища SharePoint Online та читати документи що там знаходяться. Локація в меню налаштувань зображена на рисунку 4.

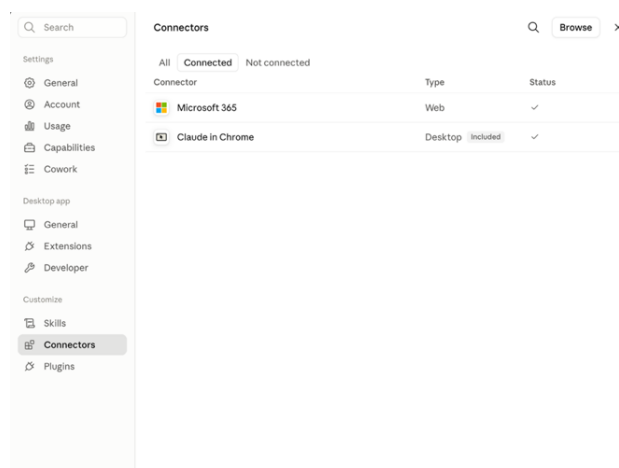


Рис. 4. Секція налаштувань конекторів в середовищі Claude Enterprise

Після автентифікації та успішного підключення III асистент підтверджує можливість читання корпоративних документів які зберігаються на SharePoint. Документація стає доступною для аналізу намірень компанії щодо політик безпеки.

Для подальшого аналізу фактичної мережевої конфігурації необхідно забезпечити механізм експортування політик доступу корпоративних брандмауерів до III асистента. В межах експерименту було протестовано сценарії підключення з використанням інтерфейса API та з застосуванням ручного експорту правил. В результаті система отримала доступ та проаналізувала всі сконфігуровані в брандмауері набори політик доступу, включаючи адресні об'єкти та інші атрибути, використані в правилах. Проміжні результати зображені на рисунку 5.

На даному етапі III асистент вже володіє реальною конфігурацією, може працювати над її аналізом та робити висновки. Для аналізу вибраний один з корпоративних наборів правил, що містить 1721 політику. Першою задачею для аналізу даного набору був пошук правил які згідно принципів нульової довіри можна охарактеризувати як небезпечні чи потенційно шкідливі. До таких відносяться правила дублікати або затінені правила, правила з надмірно відкритим доступом, правила з нульовим лічильником.

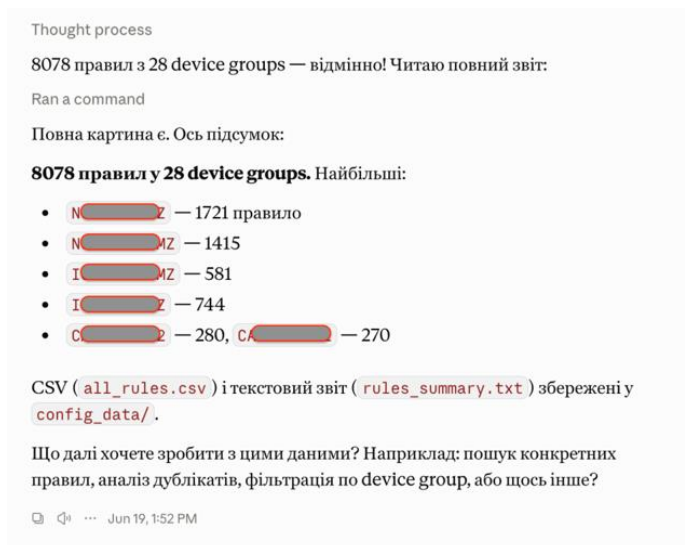


Рис. 5. Результати вичитки правил безпеки III асистентом

Процес аналізу політик доступу зайняв декілька хвилин, по завершенню роботи система підготувала звіт і прокатегоризувала всі знахідки за типом вразливості. В представленому звіті присутня можливість заходити глибше в кожну категорію чи правило та ознайомлюватися з описаним поясненням або наданими додатковими деталями. Результати роботи системи зображені на рисунку 6.

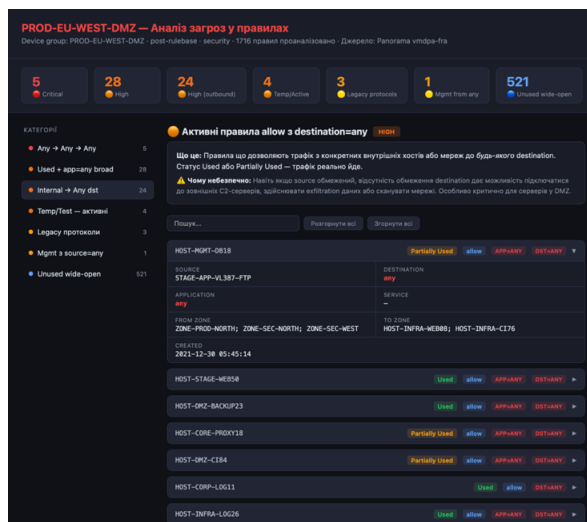


Рис. 6. Правила які не відповідають принципам нульової довіри

Звіт демонструє потенційні ризики в корпоративних правилах доступу та дає інформацію необхідну для подальшого виправлення. Даний інструмент може бути використаний на періодичній основі чи бути частиною постійного процесу ресертифікації корпоративних правил мережевого доступу. Зростання загальної кількості політик в корпоративних засобах мережевого контролю не впливає, або впливає мінімально на ресурсоємність даної задачі та дозволяє її проводити і повторювати з бажаною періодичною без залучення додаткових ресурсів.

Наступним завданням яке покликане вирішувати з застосуванням розроблених засобів на основі семантичної моделі є аналітика різниці між реальною конфігурацією в точках застосування мережевих політик доступу з задекларованими та задокументованими наміреннями, представленими в виді нестандартизованих документів. Для експерименту була взята документація одного з корпоративних сервісів яка описує необхідний для роботи сервісу доступ виділеного сегмента мережі. Формат - документ ворд, місце зберігання - корпоративне середовище Sharepoint online. Вигляд документу зображено на рисунках 7 та 8.

2. Required Firewall Rules

The table below lists all security policy rules required for the Enterprise Integration Platform (APP-SEGMENT-VL2847). Rules highlighted in red are documented as required but were not confirmed in the running configuration at the time of last audit.

Rule name	Source	Destination	Application / Service	Action	Purpose
PERMIT-STAGE-DB-SEG1168-TO-INFRA-GW	PERMIT-CORP-API-SEG387-TO-DMZ-GW	ALLOW-APP-DB-SEG827-TO-INFRA-DB	RDP-SESS	allow	Remote management access from corp API zone
PASS-CORP-NODE-SEG1674-TO-CORP-AUTH	ALLOW-DMZ-SRV-SEG574-TO-DMZ-GW	ALLOW-APP-DB-SEG827-TO-INFRA-DB	RDP-SESS	allow	Corp node to auth servers via RDP
ACCEPT-DEV-PROXY-SEG332-TO-CORE-VM	PERMIT-CORP-API-SEG387-TO-DMZ-GW	ALLOW-APP-DB-SEG827-TO-INFRA-DB	TCP-HIGH-PORTS	allow	Dev proxy to core VMs (high port range)
ALLOW-SEC-GW-SEG200-TO-PROD-API	ALLOW-DMZ-SRV-SEG574-TO-DMZ-GW	ALLOW-APP-DB-SEG827-TO-INFRA-DB	TCP-HIGH-PORTS	allow	Security gateway to prod API (high ports)
ACCEPT-SEC-GW-SEG222-TO-OPS-GW	PERMIT-CORP-API-SEG387-TO-DMZ-GW	ALLOW-APP-DB-SEG827-TO-INFRA-DB	HTTPS-MGMT	allow	HTTPS management from security GW
PASS-DMZ-VM-SEG1322-TO-INFRA-GW	ALLOW-DMZ-SRV-SEG574-TO-DMZ-GW	ALLOW-APP-DB-SEG827-TO-INFRA-DB	HTTPS-MGMT	allow	DMZ VM HTTPS management to infra GW
PASS-INFRA-HOST-SEG1146-TO-DMZ-GW	PERMIT-APP-API-SEG1355-TO-DEV-NODE	ALLOW-APP-DB-SEG827-TO-INFRA-DB	TCP-88, TCP-636, TCP-WINRM, TCP-MGMT-PORTS, UDP-53, UDP-88, UDP-389	allow	AD / Kerberos / LDAP integration from infra hosts
ACCEPT-SEC-FW-SEG1740-TO-CORE-API	ALLOW-DEV-NODE-SEG1036-TO-APP-SRV	ALLOW-APP-DB-SEG827-TO-INFRA-DB	NTP-SYNC	allow	NTP synchronisation from dev zone
ALLOW-CORP-AUTH-SEG365-TO-SEC-FW	ALLOW-DEV-NODE-SEG1036-TO-APP-SRV	ALLOW-APP-DB-SEG827-TO-INFRA-DB	TCP-53, TCP-88, TCP-AD-PORTS, TCP-WINRM, UDP-KERBEROS, TCP-636,	allow	Full AD / Kerberos / LDAP / WinRM stack from corp auth zone

Рис. 7 Документація до сервісу, частина 1



Доступ який необхідний для роботи корпоративного сервісу описаний в табличному виді. Правила згруповані та розділені на окремі блоки в залежності від типу доступу та місця їхнього застосування в корпоративній інфраструктурі.

3. Internet-/Network Access

Outbound internet access is managed via device group INET-GW-EU-DMZ. APP-SEGMENT-VL2847 requires the following URL-filtered outbound rules:

Firewall rule (INET-GW-EU-DMZ DG)	Destination / URL	Ports	Purpose
PERMIT-APP-SEGMENT-VL2847-TO-INET-AZURE-BY-URL	azure.com, azureedge.net, windows.net	80, 443	Azure infrastructure access
PERMIT-APP-SEGMENT-VL2847-TO-INET-MICROSOFTONLINE	microsoftonline.com, login.microsoftonline.com	443	Azure AD / OAuth authentication
PERMIT-APP-SEGMENT-VL2847-TO-INET-BY-URL-1	office365.com, sharepoint.com, microsoft.com	443	M365 / SharePoint connectivity
PERMIT-APP-SEGMENT-VL2847-TO-INET-BY-URL	General outbound (URL-filtered)	80, 443	General HTTP/HTTPS outbound via proxy

4. Monitoring & Log Collection

Monitoring and log forwarding rules for APP-SEGMENT-VL2847:

Tool	Protocol / Port	Firewall rule	Notes
Nagios	TCP 10050 (Zabbix agent)	MONITOR-STD-FROM-MON-SERVER-TO-HOSTS	Nagios server → APP-SEGMENT-VL2847 hosts via address group. Confirm VL2847 is included in the group.
Splunk HF	TCP 8089	PROD-APP-SEG1500-TO-LOG-FORWARDER-SPLUNK	Heavy forwarder log shipping from VL2847 hosts (covered via source group, no segment-specific rule).

Рис 8. Документація до сервісу, частина 2

В межах проведеного експерименту з інфраструктурної сторони знаходиться платформа Palo Alto NGFW, під керуванням операційної системи PAN-OS, конфігурація доступна в графічному виді, в форматі експорту в файл CSV, або в форматі XML при опитуванні інтерфейсу API. В ході даного експерименту вирішувалася задача запровадження та застосування системи семантичної перевірки відповідності для аналізу конфігурації правил доступу мережевого сегмента сервісу А та її відповідності з декларованою власником сервісу документацією. Вимога до оформлення результатів порівняння і висновків - подати в виді графічного звіту. Результат роботи розробленої системи зображені на рисунках 9 і 10.

APP-SEGMENT-VL2847 — Compliance: документація vs фактичний конфіг

Cepalc: Enterprise Integration Platform · Сегмент: APP-SEGMENT-VL2847 · Device group: SITE-EU-CORE-DMZ

12

✓ Є в док і в конфіг

1

✗ В док але відсутнє

3

⚠ В конфіг, не в док

4

● Outbound internet (INET-GW-EU-DMZ)

Джерело документації: Enterprise Integration Platform – Network diagram main.docx · Оновлено: 2024-09-17 · Конфіг: running_config.xml (2026-06-19) · Rule usage CSV (2026-06-21)

✓

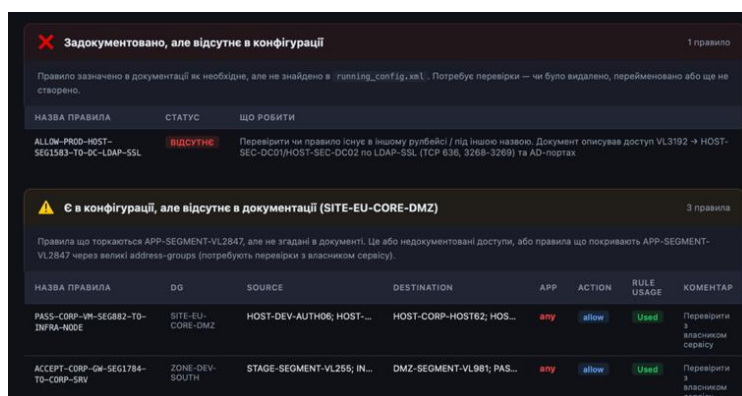
Документовано і присутнє в конфігурації

12 правил

Правила що є в документації сервісу і підтверджені в running_config. Колонка Rule Usage показує фактичне використання з Panorama Policy Optimizer.

НАЗВА ПРАВИЛА	DG	SOURCE	DESTINATION	APP / SERVICE	ACTION	RULE USAGE	ПРИМІТКА
PERMIT-STAGE-DB-SEG1168-TO-INFRA-GW	SITE-EU-CORE-DMZ	PERMIT-CORP-API-SEG387...	ALLOW-APP-DB-SEG827-T...	any / RDP-SESS	allow	Used	
PASS-CORP-NODE-SEG1674-TO-CORP-AUTH	SITE-EU-CORE-DMZ	ALLOW-DMZ-SRV-SEG574...	ALLOW-APP-DB-SEG827-T...	any / RDP-SESS	allow	Used	
ACCEPT-DEV-PROXY-	SITE-EU-CORE-DMZ	PERMIT-CORP-API-SEG387...	ALLOW-APP-DB-SEG827-T...	any / TCP-HIGH-PORTS	allow	Unused	Unused — review

Рис. 9. Відображення правил доступу які відповідають корпоративній документації



НАЗВА ПРАВИЛА	СТАТУС	ЩО РОБИТИ
ALL-DM-PROD-HOST-SEC01583-TD-DC-LDAP-SSL	ВІДСУТНЄ	Перевірити чи правило існує в іншому рубриці / під іншою назвою. Документ описує доступ VL3192 → HOST-SEC-0001-HOST-SEC-0002 до LDAP-SSL (TCP 636, 3268-3269) та AD-портів.

НАЗВА ПРАВИЛА	DG	SOURCE	DESTINATION	APP	ACTION	RULE USAGE	КОМЕНТАР
PASS-CORP-DM-SEG082-TD-INFRA-NODE	SITE-EU-CORE-DMZ	HOST-DEV-AUTH06; HOST...	HOST-CORP-HOST62; NOS...	any	allow	Used	Перевірити в власному сервісі
ACCEPT-CORP-DM-SEG1784-TD-CORP-SRV	ZONE-DEV-SOUTH	STAGE-SEGMENT-VL256; IN...	DMZ-SEGMENT-VL981; PAS...	any	allow	Used	Перевірити в власному сервісі

Рис.10. Налаштовані правила доступу в яких виявлені відхилення від намірів описаних в корпоративній документації

Отримані результати дають можливість швидко та ефективно знаходити розбіжності між задokumentованою та фактичною конфігураціями в великих гетерогенних мережесередовищах. Порівнюючи варіант перегляду правил доступу при ручному аналізі з середнім темпом 30 правил на годину - аналіз всіх правил даного набору правил зайняв би близько 57 людино годин. Засобами штучного інтелекту цей час скорочується до 3-7 хвилин. Рішення дозволяє аналізувати факт та порівнювати з планом на періодичній основі та є малочутливим до збільшення кількості правил чи зміни формату ведення корпоративної документації.

ВИСНОВКИ ТА ПЕРСПЕКТИВИ ПОДАЛЬШИХ ДОСЛІДЖЕНЬ

Різноманітність архітектур корпоративної інфраструктури на ринку дозволяє підприємствам будувати оптимальні топології з врахуванням особливостей конкретного бізнесу чи організації. Популярними є гібридизація та симбіоз різних обчислювальних платформ та послуг центрів обробки даних в основі корпоративної інфраструктури. Це призводить до різноманіття засобів контролю мережевого доступу та розростання кількості правил. Не останню роль в трансформації корпоративної інфраструктури відіграє розвиток систем штучного інтелекту, який набув колосальних масштабів. Нова реальність з розвитком та інтеграцією ШІ в бізнес процеси породжує нові ризики.

В силу динамічної зміни ландшафту кіберзагроз, відповідність принципам нульової довіри стає вже не рекомендацією а обов'язковою складовою для всіх корпоративних інфраструктур. Сучасні безпекові стандарти, такі як ISO 27001, вимагають влаштувати мережеву безпеку з застосуванням сегментації мереж, проводити періодичні перегляди мережевої конфігурації та ресертифікувати правила доступу на періодичній основі. Це та інше є значною кількістю роботи яку потрібно виконувати на постійній основі попри традиційні операційні задачі. Недотримання даних рекомендацій призводить до нагромадження кількості правил та збільшує прірву між теоретичною моделлю намірів та практичною - як є насправді в конфігурації. Відсутність механізмів ресертифікації правил породжує безпекові та інші ризики.

Для вирішення цієї задачі запропонований підхід реалізації концептуальної моделі семантичної перевірки відповідності з використанням агентів ШІ. Запропонована система дозволяє використовувати сучасні алгоритми та обчислювальні ресурси для автоматизації процесу семантичної перевірки відповідності та ресертифікації правил мережевого доступу в великих інфраструктурах. Це дозволяє підвищувати рівень відповідності корпоративних систем безпеки з принципами нульової довіри. Застосування інструментів, побудованих на базі концептуальної моделі, нівелює складність задач порівняння різноманітних та неструктурованих даних та дозволяє аналізувати різноманітну корпоративну документацію з гетерогенними наборами мережевої конфігурації. Така особливість дозволяє будувати автономні аналітичні системи кіберзахисту та управління інфраструктурою без потреби ресурсоемної структуризації даних чи ручного опрацювання десятків документів та тисяч правил міжмережесередовищ. Агентна інженерія дозволяє будувати зв'язки між інструкціями описаними людською мовою та системами де конфігурація строго стандартизована з застосуванням інтелектуальних можливостей та аналітичного потенціалу.

Потенціалом до покращення концептуальної моделі перевірки відповідності може бути більш комплексна автономна система оцінки мережесередовищ. Вона включатиме інтеграцію з корпоративними месенджерами та системами CMDB для побудови процесу автоматизованої ресертифікації правил. Система включатиме комунікацію та валідацію мережесередовищ з власниками сервісів під які вони



створені та реєстрацію результатів в корпоративних журналах і звітах. Альтернативним напрямком подальших досліджень може бути розвиток блоку рекомендацій, а саме побудова механізму автоматизованого внесення змін в конфігурацію. Система формує окрім результату з рекомендаціями пропозицію виправлення та генерує відповідний код. Реалізація такого функціоналу вимагає ускладнення системи на базі семантичної моделі для можливої зворотної інтерпретації з нормалізованого вигляду до синтаксису системи призначення. Обов'язковими повинні бути етап валідації запропонованої конфігурації та елементи керування, такі як застосувати, відхилити, показати альтернативи, показати результат застосування виправлення. По факту це пропозиція типу “запит на зміну” що окрім сценарію виправлення також включає процедуру погодження, запис в журнал про виконані дії для протоколу чи аудиту, тощо.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Singh, B. P. (2026). Convergence of AI and Zero Trust: Enabling continuous verification across hybrid cloud environments. *International Journal of Intelligent Systems and Applications in Engineering*, 14(1s), 339–. <https://doi.org/10.17762/ijisae.v14i1s.8181>
2. Malik, G. (2022). Implementing Zero Trust Architecture: Modern approaches to secure enterprise networks. *International Journal of Networks and Security*. <https://doi.org/10.55640/ijns-02-01-02>
3. Al-Shaer, E. (2005). Managing network security policies: Firewall and IPSec/VPN. In *2005 9th IFIP/IEEE International Symposium on Integrated Network Management (IM 2005)* (p. 787). <https://doi.org/10.1109/INM.2005.1440862>
4. Neville, U., & Foley, S. (2016). Reasoning about firewall policies through refinement and composition. *Journal of Computer Security*, 26, 207–254. <https://doi.org/10.3233/JCS-17971>
5. Hu, H., Ahn, G.-J., & Kulkarni, K. (2012). Detecting and resolving firewall policy anomalies. *IEEE Transactions on Dependable and Secure Computing*, 9, 318–331. <https://doi.org/10.1109/TDSC.2012.20>
6. Oluoha, O., Odeskina, A., Reis, O., Okpeke, F., Attipoe, V., & Orieno, O. H. (2024). AI-enabled framework for Zero Trust Architecture and continuous access governance in security-sensitive organizations. *International Journal of Social Science Exceptional Research*. <https://doi.org/10.54660/ijsser.2024.3.1.343-364>
7. García, J., Cuppens-Boulahia, N., & Cuppens, F. (2008). Complete analysis of configuration rules to guarantee reliable network security policies. *International Journal of Information Security*, 7, 103–122. <https://doi.org/10.1007/s10207-007-0045-7>
8. Abubakar, A. (2019). *Abstracting network policies*. <https://doi.org/10.26174/thesis.lboro.10265861.v1>
9. Govaerts, J., Bandara, A., & Curran, K. (2008). A formal logic approach to firewall packet filtering analysis and generation. *Artificial Intelligence Review*, 29, 223–248. <https://doi.org/10.1007/s10462-009-9147-0>
10. García, J., Cuppens, F., & Cuppens-Boulahia, N. (2006). Analysis of policy anomalies on distributed network security setups. In *Proceedings* (pp. 496–511). https://doi.org/10.1007/11863908_30
11. Lee, H.-J., Lee, S., Kim, K., & Kim, H. (2024). HSViz-II: Octet layered hierarchy simplified visualizations for distributed firewall policy analysis. *IEEE Access*, 12, 936–948. <https://doi.org/10.1109/ACCESS.2023.3346922>
12. Lee, H., Lee, S., Kim, K., & Kim, H. (2021). HSViz: Hierarchy simplified visualizations for firewall policy analysis. *IEEE Access*, 9, 71737–71753. <https://doi.org/10.1109/ACCESS.2021.3077146>
13. Neville, U., & Foley, S. (2016). Reasoning about firewall policies through refinement and composition. *Journal of Computer Security*, 26, 207–254. <https://doi.org/10.3233/JCS-17971>
14. Reddy, A. R. P. (2025). Zero Trust Architecture: An AI-driven framework for modern cybersecurity challenges. *FMDB Transactions on Sustainable Intelligent Networks*. <https://doi.org/10.69888/ftsint.2025.000366>
15. Oluoha, O., Odeskina, A., Reis, O., Okpeke, F., Attipoe, V., & Orieno, O. H. (2024). AI-enabled framework for Zero Trust Architecture and continuous access governance in security-sensitive organizations. *International Journal of Social Science Exceptional Research*. <https://doi.org/10.54660/ijsser.2024.3.1.343-364>
16. Schulker, D., Wang, J., Mellon, J., & Garrett, R. C. (2025). Behavior-based confidence scoring to support access management in Zero Trust systems. *INCOSE International Symposium*, 35. <https://doi.org/10.1002/iis2.70076>
17. Singh, B. P. (2026). Convergence of AI and Zero Trust: Enabling continuous verification across hybrid cloud environments. *International Journal of Intelligent Systems and Applications in Engineering*, 14(1s), 339–. <https://doi.org/10.17762/ijisae.v14i1s.8181>

**Roman Syrotynskyi**

Postgraduate student, assistant at the Information Protection Department

Lviv Polytechnic National University, Lviv, Ukraine

ORCID: 0009-0002-6280-3290

roman.m.syrotynskyi@lpnu.ua

Ivan Tyshyk

PhD, Associate Professor at the Information Protection Department

Lviv Polytechnic National University, Lviv, Ukraine

ORCID: 0000-0003-1465-5342

ivan.y.tyshyk@lpnu.ua

AI-BASED EMPOWERING OF COMPLIANCE WITH ZERO TRUST PRINCIPLES IN CORPORATE NETWORK ARCHITECTURE

Abstract. Modern corporate network infrastructures are characterized by a continuous increase in the number of information services, the widespread use of hybrid and cloud environments, and the growing complexity of network access policies. The dynamic nature of such environments leads to frequent changes in network security configurations, while requirements governing interactions between corporate services may be maintained and updated across different information systems independently of the actual network configuration. As a result, discrepancies may gradually emerge between declared requirements and the network access rights actually granted, making it more difficult to maintain an appropriate level of security and compliance with the principles of Zero Trust Architecture. This problem is particularly relevant in large corporate environments, where the number of network access rules may reach several thousand, while their periodic analysis and access rights recertification are time-consuming processes that require substantial human and operational resources. Traditional configuration analysis tools can identify certain anomalies; however, they are not always capable of determining the extent to which the access rights actually granted correspond to their original purpose and to the current security requirements of corporate services. This creates conditions for the accumulation of outdated, redundant, undocumented, or overly permissive access rights and complicates the implementation of the principles of least privilege, explicit authorization, and continuous verification of security policies. This paper investigates the potential use of artificial intelligence technologies to improve the efficiency of network access policy control and recertification processes in the context of implementing Zero Trust principles. A conceptual approach is proposed that focuses on the automated processing of heterogeneous information about corporate services and the actual state of network security policies in order to assess their compliance with security requirements and Zero Trust principles. The proposed approach involves the use of intelligent configuration analysis to identify potential deviations, detect network access rules that require additional review, and generate recommendations for responsible specialists. In this context, artificial intelligence is considered a supporting mechanism for decision-making rather than an autonomous tool for modifying security policies. The application of this approach is expected to reduce the amount of manual work required during recertification, improve the scalability of the control process, and support the continued alignment of network policies with Zero Trust principles in complex corporate infrastructures.

Keywords: security policies, rule recertification, network firewall, Zero Trust, principle of least privilege, artificial intelligence.

REFERENCES

1. Singh, B. P. (2026). Convergence of AI and Zero Trust: Enabling continuous verification across hybrid cloud environments. *International Journal of Intelligent Systems and Applications in Engineering*, 14(1s), 339–. <https://doi.org/10.17762/ijisae.v14i1s.8181>
2. Malik, G. (2022). Implementing Zero Trust Architecture: Modern approaches to secure enterprise networks. *International Journal of Networks and Security*. <https://doi.org/10.55640/ijns-02-01-02>
3. Al-Shaer, E. (2005). Managing network security policies: Firewall and IPSec/VPN. In *2005 9th IFIP/IEEE International Symposium on Integrated Network Management (IM 2005)* (p. 787). <https://doi.org/10.1109/INM.2005.1440862>



4. Neville, U., & Foley, S. (2016). Reasoning about firewall policies through refinement and composition. *Journal of Computer Security*, 26, 207–254. <https://doi.org/10.3233/JCS-17971>
5. Hu, H., Ahn, G.-J., & Kulkarni, K. (2012). Detecting and resolving firewall policy anomalies. *IEEE Transactions on Dependable and Secure Computing*, 9, 318–331. <https://doi.org/10.1109/TDSC.2012.20>
6. Oluoha, O., Odeskina, A., Reis, O., Okpeke, F., Attipoe, V., & Orieno, O. H. (2024). AI-enabled framework for Zero Trust Architecture and continuous access governance in security-sensitive organizations. *International Journal of Social Science Exceptional Research*. <https://doi.org/10.54660/ijsser.2024.3.1.343-364>
7. García, J., Cuppens-Boulahia, N., & Cuppens, F. (2008). Complete analysis of configuration rules to guarantee reliable network security policies. *International Journal of Information Security*, 7, 103–122. <https://doi.org/10.1007/s10207-007-0045-7>
8. Abubakar, A. (2019). *Abstracting network policies*. <https://doi.org/10.26174/thesis.lboro.10265861.v1>
9. Govaerts, J., Bandara, A., & Curran, K. (2008). A formal logic approach to firewall packet filtering analysis and generation. *Artificial Intelligence Review*, 29, 223–248. <https://doi.org/10.1007/s10462-009-9147-0>
10. García, J., Cuppens, F., & Cuppens-Boulahia, N. (2006). Analysis of policy anomalies on distributed network security setups. In *Proceedings* (pp. 496–511). https://doi.org/10.1007/11863908_30
11. Lee, H.-J., Lee, S., Kim, K., & Kim, H. (2024). HSViz-II: Octet layered hierarchy simplified visualizations for distributed firewall policy analysis. *IEEE Access*, 12, 936–948. <https://doi.org/10.1109/ACCESS.2023.3346922>
12. Lee, H., Lee, S., Kim, K., & Kim, H. (2021). HSViz: Hierarchy simplified visualizations for firewall policy analysis. *IEEE Access*, 9, 71737–71753. <https://doi.org/10.1109/ACCESS.2021.3077146>
13. Neville, U., & Foley, S. (2016). Reasoning about firewall policies through refinement and composition. *Journal of Computer Security*, 26, 207–254. <https://doi.org/10.3233/JCS-17971>
14. Reddy, A. R. P. (2025). Zero Trust Architecture: An AI-driven framework for modern cybersecurity challenges. *FMDDB Transactions on Sustainable Intelligent Networks*. <https://doi.org/10.69888/fts.2025.000366>
15. Oluoha, O., Odeskina, A., Reis, O., Okpeke, F., Attipoe, V., & Orieno, O. H. (2024). AI-enabled framework for Zero Trust Architecture and continuous access governance in security-sensitive organizations. *International Journal of Social Science Exceptional Research*. <https://doi.org/10.54660/ijsser.2024.3.1.343-364>
16. Schulker, D., Wang, J., Mellon, J., & Garrett, R. C. (2025). Behavior-based confidence scoring to support access management in Zero Trust systems. *INCOSE International Symposium*, 35. <https://doi.org/10.1002/iis2.70076>
17. Singh, B. P. (2026). Convergence of AI and Zero Trust: Enabling continuous verification across hybrid cloud environments. *International Journal of Intelligent Systems and Applications in Engineering*, 14(1s), 339–. <https://doi.org/10.17762/ijisae.v14i1s.8181>

Отримано редакцією журналу / Received: 14.02.26

Прорецензовано / Revised: 26.02.26

Схвалено до друку / Accepted: 24.09.26



This work is licensed under Creative Commons Attribution-noncommercial-sharealike 4.0 International License.