

АНАЛІЗ СТІЙКОСТІ СТЕГАНОАЛГОРИТМІВ ДО ГЕОМЕТРИЧНИХ АТАК ТА СТИСНЕННЯ

Проведено аналіз стійкості стеганографічних методів приховування в просторовій області до поширених у відкритих та закритих каналах атак, таких як геометричні перетворення та стиснення. Оцінено вплив цих атак на стабільність роботи стеганоканалу та запропоновано методи підвищення стійкості до таких атак. Досліджено різні типи атак на основі афінних перетворень та їх вплив на якість приховування інформації. Зосереджено увагу на збереженні цілісності повідомлення під час таких атак.

Дослідження оцінює можливість використання різних типів повідомлень та вплив кодування з повторенням на стійкість стеганосистеми. Розроблено та змодельовано метод оцінювання стійкості атак на основі методу бітової похибки, який був модифікований для проведення аналізу на байтовому рівні задля застосовуваності результатів до різних типів контейнерів та підтримки потокових даних.

Зазначено важливість відсутності впливу порядку розташування пікселів при вбудовуванні повідомлення, що дозволяє мінімізувати вплив геометричних перетворень на вбудоване повідомлення. Стійкість алгоритму модифікації кольорового простору до атак вища, ніж в методу LSB на 14-98%. LSB має вищу на 10% стійкість лише до атак незначного масштабування. Проте ця різниця не дозволяє назвати його практично придатним до використання.

Проаналізовано вплив роздільної здатності зображення на обсяг вбудованої інформації в кожному з методів. Проаналізовано перспективні напрями проведення подальших досліджень з допомогою введення блокового кодування і використання ефективних методів корекції помилок та використання альтернативних властивостей контейнера для вбудовування.

Ключові слова: стеганографія, геометричні атаки, афінні перетворення, кольорний простір.

Вступ

Тенденція стрімкого зростання обсягу цифрової інформації, а також частоти щоденної комунікації постійно збільшує потребу в надійному та безпечному обміні приватною або комерційною інформацією. Через таке збільшення обсягу інформації, що передається, в сьогоdnішньому світі питання захисту даних від неавтентифікованого та не авторизованого доступу стає все більш актуальним. Звичайного шифрування не достатньо для гарантування безпеки інформації, часто для її захисту залучається стеганографія, яка забезпечує не лише приховування даних, але й сам факт передачі. Доволі часто за стеганоконтейнер використовують цифрові зображення. Завдяки поширенню та збільшенню доступності більш потужних відеоприскорювачів стеганографія у 3D просторі, а також у віртуальному просторі (VR) набуває все більшої популярності. Зберігається потреба шифрувати дані та зростає потреба їх приховувати. Це завдання вирішує стеганографія, особливо з використанням зображень (текстур), оскільки вони мають більшу пропускну здатність та є поширеними на сьогодні способами вбудовування повідомлення, цифрового підпису чи водяного знаку.

Формулювання проблеми

Для підвищення безпеки та стійкості стеганоканалу до атак виникає потреба дослідження стійкості популярних методів до ймовірних атак. В даній роботі проведено дослідження рівня стійкості найуживанішого методу приховування в найменш значущий біт, а також нетипового методу з використанням кольорового простору до найпоширеніших геометричних атак та стиснення, у яких приховано стеганоповідомлення.

Мета та завдання дослідження

Метою роботи є дослідити методи підвищення стійкості до конкретних атак та оцінити їх вплив на стабільність роботи стеганоканалу. Для досягнення зазначеної мети визначено такі основні завдання дослідження:

1. Розробити метод оцінки стійкості стеганографічних алгоритмів інваріантних до поширених атак;
2. Провести моделювання атак та порівняння стійкості стеганографічних алгоритмів;
3. Проаналізувати засоби покращення стійкості методів приховування.

Аналіз літератури

Найпоширенішими атаками на зображення у відкритих та закритих каналах передачі даних є трансформації, такі як афінні перетворення [1], та додавання шуму. Переважно ці методи дозволяють видозмінити зображення, проте ці зміни відразу помітні.

До прихованих методів атаки належить: незначне стиснення; вирізання ряду пікселів незначної ширини; афінне перетворення задля масштабування геометричних розмірів зображення на незначний розмір [1]; незначна кольорокорекція, що призводить до перерахунку значень інтенсивностей усіх пікселів. Завдання цих атак – здійснити зміни, які є непомітними для ока людини, залишити зображення придатним до використання та модифікувати чи унеможливити зчитування повідомлення зі стеганоконтейнера.

У наукових працях поширеним підходом до аналізу ефективності при розробці нових методів є використання метрик, які не забезпечують комплексного аналізу спотворень стеганографічного повідомлення. Найбільш популярними метриками залишаються показники: співвідношення сигналу до шуму (SNR), пікового співвідношення сигналу до шуму (PSNR), індексу структурної подібності (SSIM) [2], а також показники місткості вбудовування. Рідко проводиться порівняння із поширеними методами вбудовування. Одним з аспектів, якому не приділяють достатньо уваги, є стійкість розроблених методів до типових атак. Некоректно стверджувати, що цим питанням не приділяється достатньо уваги, проте, у більшості випадків, аналіз обмежується вивченням лише окремих типів атак, як у дослідженні Zhang Y. [3] або специфічних методів виявлення. На противагу цьому, методи виявлення демонструють значно вищий рівень системності, наприклад у роботі Arau R. [4]. Розробка цих методів, здійснюється на основі чітко визначених, хоча й обмежених за обсягом, тестових наборів даних. Це сприяє як повторюваності результатів, так і об'єктивному порівнянню ефективності.

Прикладом такого дослідження є метод [5], який реалізує комплексний аналіз впливу атак на можливість виявлення. Ця робота розширює зазначене дослідження через врахування параметра стійкості методів під впливом атак. Це дозволить знайти баланс між стійкістю і прихованістю, адже обидві характеристики важливі для практичного застосування.

Досліджені стеганографічні алгоритми

Досліджено один із найпоширеніших алгоритмів приховування у просторовій області метод модифікації найменш значущого біта (Least Significant Bit). Основний принцип роботи цього методу ґрунтується на зміні найменш значущого біта з метою приховування даних. Цей біт має незначний вплив, тому такі зміни є непомітними. Метод є незахищеним від багатьох типів атак, адже мінімальні зміни в бітових послідовностях можуть спричинити появу шуму, що виявляється при аналізі [6].

Сучасний підхід у стеганографії включає використання методів, що модифікують колірний простір. Це дає можливість приховувати інформацію не в окремих пікселях, а в сукупності кольорів, котрі використовуються в зображенні. Наразі ці методи набувають популярності, адже сучасні методи статистичного аналізу зазвичай не вивчають множину кольорів. Зміни в колірному просторі мають рівномірний вплив на градієнти, а також на статистичні характеристики контейнера. Колірний простір вважається менш вразливим до атак, які стискають або змінюють зображення, оскільки значна частина таких атак не впливає на колір. Про високу стійкість до вирізання ряду пікселів зображення свідчить те, що втрата пікселів не обов'язково вплине на множину кольорів, завдяки чому вона не зменшується [7]. Розглянуто адаптивний метод зміни колірного простору, запропонований Margalikas E. та Ramanauskaitė S. [7]. Цей метод використовує RGB-куб для кодування повідомлень, у якому кольори зображення розміщені в тривимірному просторі. Кожен з кольорів задається координатами, котрі обчислюються на основі значень пікселів (R, G, B).

Процес вбудовування повідомлення передбачає рекурсивне розбиття RGB-куба на менші підкуби, поки кожен підкуб містить більше одного кольору (рис. 1). Якщо розмір ребра

куба більше одиниці, підкуби з одним кольором можна використовувати для вбудовування даних, змінюючи значення однієї або декількох координат. Кількість вбудованих даних залежить від розмірів підкуба.

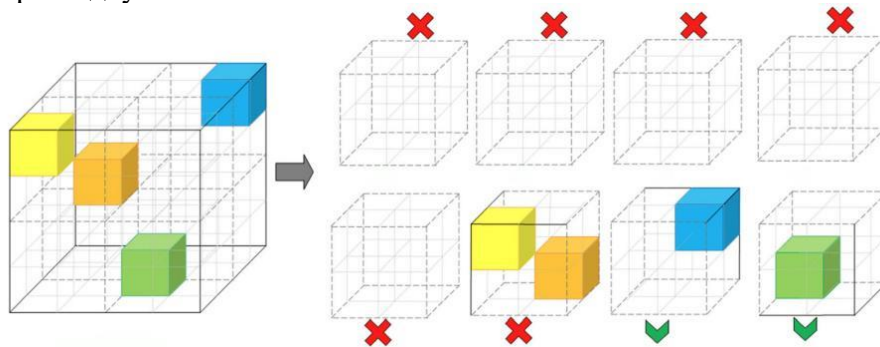


Рис 1. Відбір підкубів, у які буде здійснено приховування даних

Цей метод має недолік у вигляді неконтрольованої рекурсивності, що може призвести до спотворень і збільшення дисперсії в кольорному просторі, а це спричиняє значні статистичні й візуальні зміни. Задля мінімізації цього недоліку використано модифікацію, що враховує особливості людського ока. Для цього визначено коефіцієнти граничних змін кольорів (1), щоб виконати всі зміни в просторі синього кольору. Також обмежено дельту змін (2) до одиниці, визначивши максимальний розмір куба рівним 2×2 [8].

$$k_r = 0; k_g = 0; k_b = 1; \quad (1)$$

$$\Delta_c = 1; \quad (2)$$

Методи дослідження

Комп'ютерне моделювання досліджуваних алгоритмів вбудовування реалізовано на мові Python, з використанням бібліотек для роботи з масивами даних та векторами SciPy [9], NumPy [10], Pandas [11]. Оскільки метод модифікації кольорного простору має більшу обчислювальну складність, усі обчислення векторизовані задля уникнення циклів по всіх пікселях зображення. Геометричні атаки віддзеркалення та масштабування на один піксель [1] реалізовано за допомогою бібліотеки CV2 [12]. Для атаки стиснення використано бібліотеку Pillow [13] з налаштуваннями якості $q=0,95$, що вносить незначні спотворення до зображення та прибирає високочастотні шуми та деталі. Геометричне обрізання (Crop) виконано бібліотекою Pillow на масиві пікселів.

Бібліотека CV2 працює з обчисленнями з рухомою комою, що вносить похибку в обчислення. Для порівняння додано метод віддзеркалення через сортування масиву пікселів. Для атаки обрізання експериментальним шляхом встановлено розмір смуг обрізання по 10–20 пікселів як оптимальне число, яке відображає різницю між методами.

Моделювання методів частотного простору

У межах цього дослідження було проведено аналіз методів приховування інформації в частотній області. Для моделювання використано метод вбудовування на основі дискретно косинусного перетворення, а саме моделювання проведено на основі методів реалізованих у відкритому програмному забезпеченні. Використання DC (Direct Current – перший коефіцієнт для базової частоти, відповідає за середнє значення) коефіцієнтів для вбудовування, як запропоновано у роботі Shivam M. [14], спотворює кольорну гістограму зображення і факт вбудовування видно неозброєним оком. Такий метод був відкинутий, тому що він потребує доопрацювання.

Для моделювання вбудовування змін в найменш значущі біти АС коефіцієнтів (Alternate Current – всі коефіцієнти крім першого, що показують відхилення на відповідній частоті) використано бібліотеку SciPy, модулі fftpack та fft [15]. Обидва модулі працюють з

коефіцієнтами як з числами з рухомою комою. При побітових обчисленнях втрачається точність, що вносить похибку в повідомлення під час вбудовування, ще до застосування атак. Введення матриці повторень коефіцієнтів, для одного біту повідомлення, не дало результату і цей метод також був відкинтий.

Зважаючи на вищенаведені функціональні обмеження моделей, результати можна вважати нерепрезентативними та такими, що потребують додаткового дослідження. Його доцільно провести на основі більш значущих бітів коефіцієнтів або ж із застосуванням інших методів вбудовування в частотну область. Адже вбудовування в частотні коефіцієнти є перспективним методом для подальших досліджень впливу геометричних атак. Однак у рамках даного дослідження обрана виключно просторова область.

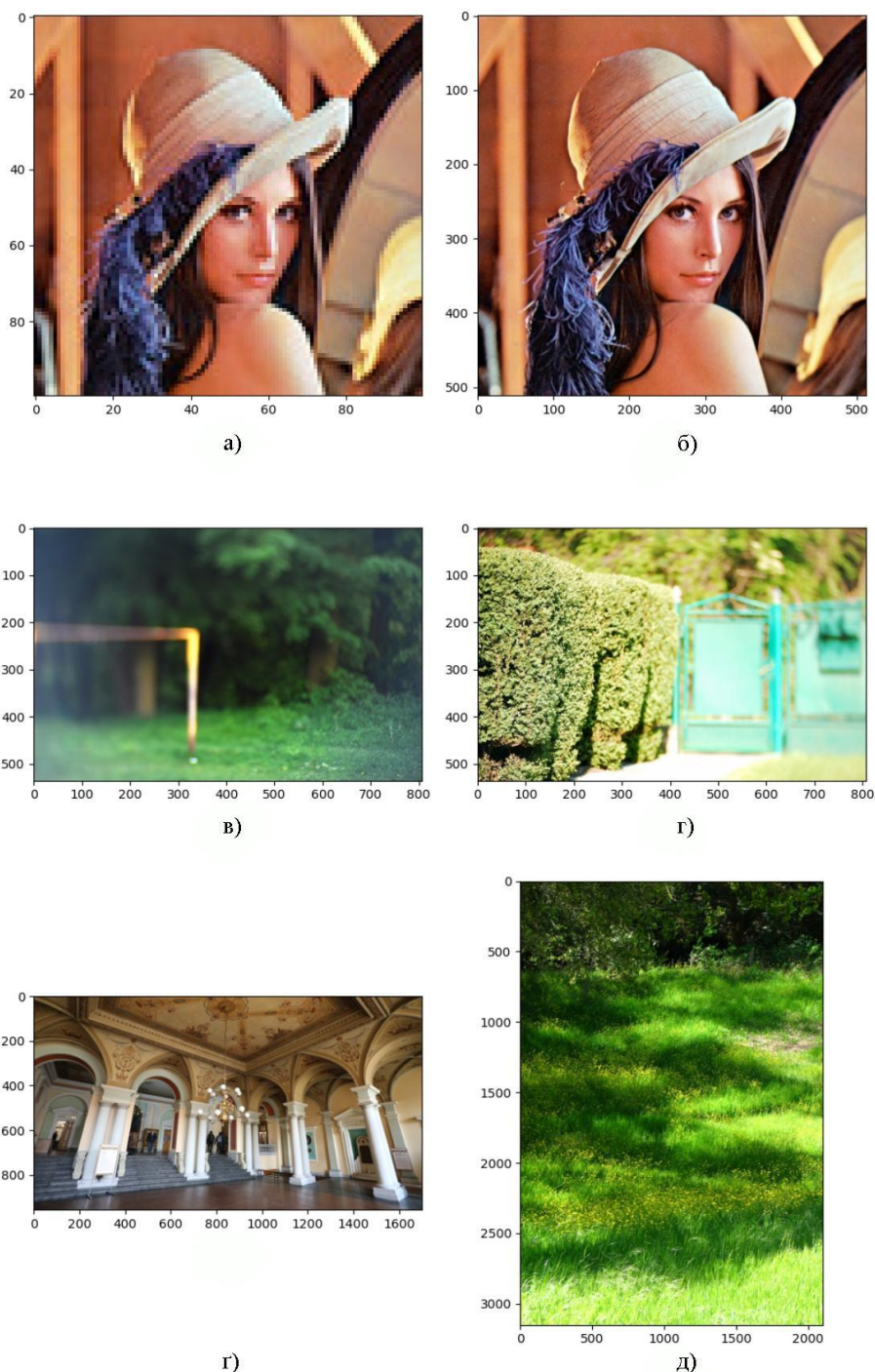


Рис 2. Основні типи використаних зображень

Тестові дані

Для покриття більшої кількості можливих вхідних даних використано декілька наборів зображень. Основні типи зображень (рис. 2) містять зображення різного розміру та різних ступенів стиску (а,б). Також підібрано зображення з різними рівнями інтенсивності блакитного каналу (в,г) та різними значеннями роздільної здатності (г,д).

Для обчислення отриманого значення стану вбудовуваного повідомлення після певної атаки бралось медіанне значення після опрацювання різних зображень. Рівень розкиду параметра стійкості зазвичай становив $\pm 15\%$, залежно від використаного зображення. У випадку сприятливих умов параметри атаки змінювались, щоб атака зачепила хоча б якийсь піксель з вбудованою інформацією. У випадках великої різниці в результатах їх згруповано у дві групи та внесено у таблиці як середнє значення першої та другої груп.

Методи оцінювання

Оскільки не всі методи вбудовування працюють з блоковою репрезентацією інформації, виникає потреба у дослідженні методу оцінювання ступеня інваріантності до атак. Згідно з науковою літературою, найчастіше для таких цілей застосовується показник бітової похибки (Bit-Error Rate, BER) [16]. Втім, на основі результатів моделювання запропоновано розробити та протестувати показник байтової похибки (Byte-Error Rate), який не піддається впливу помилково-позитивних статистичних збігів під час роботи з потоковими даними у випадку зникнення одного біту. Оцінювання ідентичне до Bit-Error Rate, проте проводиться на основі байтів. Це дозволить порівняти відсоток інформації, що підлягає видобуванню після атаки, та відображає ефективність вбудовування зображення, бінарного архіву чи цифрового ключа. До прикладу QR-код, котрий може бути використаний як зображення для водяного знаку, має істотний захист від значних змін, внесених модифікаціями. Виконувати файли з перевіркою контрольних сум недоцільно передавати або ж слід припиняти їх використання у випадку атаки на канал.

Кодування з виправленням

Під час дослідження виявлено, що для всіх методів модифікація чи втрата одного біта внаслідок атаки приводила до зсуву подальших бітів, що змушує зсунуті біти вважати помилкою. Проведено моделювання виявлення та усунення таких помилок з застосуванням кодування з виправленням, що не дало позитивного результату. Оскільки в розглянутих атаках біти частіше зникали, а не модифікувались, застосовано метод кодування з виправленням Ріда-Соломона [17]. Це не покращило результати, адже для цього методу робота на потокових даних чи на блоках великого розміру з пропущеними бітами не є ефективним режимом роботи. Задача виявлення пропущених бітів виконується на блоках меншого розміру. Зсунуті біти вважаються помилками, тож така значна їх кількість перевищує допустиму межу. Планується продовжити дослідження методів кодування з виправленням помилок та зникненням бітів із застосуванням метода Ріда-Соломона і розбиттям потокових даних на блоки.

Кодування повторенням

З огляду впливу атак на текстове повідомлення (табл. 1) прийнято рішення внести в систему кодування повторенням, яке захистить повідомлення від значних втрат через зникнення одного пікселя. Суть методу зводиться до повторення кожного біту повідомлення визначену кількістю разів n . Внаслідок повторення, один або декілька бітів у групі змінюються, інші ж залишаються незмінними, що дає можливість виявити та виправити помилку методом кворуму [18]. Це покращило наочність порівняння, адже втрата одного пікселя більше не призводить до втрати читабельності всього тексту, що йде після цього пікселя.

Цей вид кодування ефективний, коли атака вносить помилку. У випадку втрати даних, застосовано модифікацію, яка передбачає повторення кожного біта декілька разів для забезпечення відновлення даних за втрати частини бітів. При втраті під час передачі або зберіганні, надлишковість, створена повторенням, дає змогу відновити початкове

повідомлення. Зокрема, якщо "1" повторюється тричі як "111", а один біт зникає, залишаючи "11", то система може визначити, що вихідне значення "1". Цей підхід є легким у реалізації, бо не потребує складних обчислень. Кодування повторенням бітів істотно збільшує об'єм вбудованих даних, що непрактично в системах з обмеженими ресурсами і є ефективним лише для виявлення та виправлення окремих помилок у групах бітів. Слід зауважити, що ця модифікація знижує можливість виправлень помилок, адже фокусується на виявленні пропущених значень, таким чином похибку "101" може розпізнати як перехід до наступного значення [18].

Таблиця 1

Вплив застосування атаки на зміни повідомлення побайтово

Атака	Найменш значущих біт, %	Модифікації колірного простору, %
Віддзеркалення (CV2)	0.55	38-100
Масштабування на 1 піксель	7.77	0.55
Віддзеркалення (Array)	0	100
Вирізання ряду пікселів	0.55	99
Стиснення JPG (q=0.95)	0	0
Обрізання рамки (10-20px)	0	37

Результати

Після додавання кодування покращилась стійкість системи, а також ефективність видобування читабельного повідомлення у випадку роботи з текстом (табл. 3). Текст латинськими символами більш стійкий до спотворення, адже для кодування символу в ньому використано один байт в системах кодування ASCII та UTF-8. Для кирилиці це значення дорівнює двом байтам (UTF-8), а для рідкісних алфавітів та спеціальних символів може досягати чотирьох, що погіршує читабельність відновленого тексту [19]. З таблиці 2 видно, що втрата невеликої кількості біт призводить до спотворення не лише поточного, але й подальших символів, при загальній схожості бітів у 99,43% у зразку для методу модифікації колірного простору.

Результати дослідження надають змогу розширити звичне порівняння методів вбудовування за допомогою нових практичних параметрів, таких як стійкість до геометричних атак та стиснення. Такий аналіз дозволяє виконувати підбір методу на основі характеристик каналу зв'язку та можливих атак. Це розширює можливості розповсюджених порівнянь сигналу до шуму, наприклад у роботі Walia [2], та робить таке порівняння ширшим і більш практичним. Розширений набір змодельованих атак продовжує роботу Zhang Y. [3] та доповнює роботу Alrusaini O. [5]. Результати порівняння методів просторової області свідчать про доцільність збільшення кількості атак для комплексного аналізу, адже стійкість методу вбудовування різниться залежно від властивостей атаки.

Результати порівняння двох методів вбудовування (табл. 3) демонструють, що при атаці віддзеркалення суттєво кращі результати показує метод модифікації колірного простору, адже в ньому немає залежності на послідовність пікселів. Вирізання ряду пікселів та обрізання рамки по-різному впливає на метод вбудовування в найменш значущі біти залежно від того скільки пікселів вбудованого зображення знаходиться у цьому ряді, а на колірний простір

втрата одного ряду пікселів впливу майже немає. Вплив вирізання рамки прямопропорційний розміру рамки.

Таблиця 2

Приклади декодованих повідомлень

Метод	Статистична схожість байт, %	Стан повідомлення після атаки
Оригінал	100	123 Hello Alice.This is Bob. Im not sure about what limit do we have, it is probably quite big. Nevertheless message should be long enough to live through randomly disappearing bits.
Найменш значущих біт	10	123 Hello □□□,l□5CR□B□□□Z□r□□R□d□□□□UTb□G5□□)qYj*□ #□UR□□Uj□□jmj□JQN□□D□UUD□T□F□TJ□UU□□UJw□kPY:L □\$R□□□□)□H□□N□*□□J□□□%□iQUUJ□□U□□□UJ□j□□R□UJ □UTUU□□□UV□□□□□VZts□Jz□eTR□□□U*□U□&jmVV□□□□ mUYJ□□V□□I□YB□V□YUU*□*□□ДК□□□1UWJ./Ji2)'U□ѿ□□ □□#U□□□V□F
Модифікація кольорового простору	99.43	123 Hello Alice.This is Bob. Im not sure about what limit do we have, it is probably quite big. Nevertheless message should be long2□7□□□:7□64□2□:497□□□90□27□□<□24□□□2□□4□3□14□ 9

Таблиця 3

Вплив застосування атаки на зміни закодованого повідомлення побайтно

Атака	Найменш значущих біт, %	Модифікації кольорового простору, %
Віддзеркалення (CV2)	0	99.4
Масштабування на 1 піксель	10	0.74
Віддзеркалення (Array)	0.92	100
Вирізання ряду пікселів	0.55-99	99
Стиснення JPG (q=0.95)	0.55	0.55
Обрізання рамки (10-20px)	0.74	15

Для методу вбудовування в найменш значущі біти розмір зображення визначає обсяг інформації, яку можна вбудувати, та незначно впливає на стійкість до атак. Максимальний обсяг вбудованого повідомлення методом модифікації кольорового простору не значно корелює з роздільною здатністю контейнера, зокрема від кількості пікселів [14]. Натомість є залежність від дисперсії кольорового простору. Чим більше унікальних пікселів існує і чим більше вони розосереджені по кольоровому простору, тим більший обсяг інформації можна вбудувати. Стійкість до атак залежить від кількості повторень кожного унікального пікселя: чим більше повторень, тим стійкішим буде вбудовування завдяки тому, що значення в кубі кольорового простору будуть представлені кількома пікселями в зображенні. Це визначний фактор при підборі зображень-контейнерів.

Атака незначного масштабування показує незначну перевагу методу вбудовування в найменш значущий біт. При атаці стисненням перераховуються усі значення кольорів. Тож жоден метод у порівнянні не можна вважати стійким до неї.

Наукова новизна отриманих результатів дослідження полягає у тому, що розроблено метод оцінки стійкості стеганоалгоритмів з урахуванням впливу різних атак. Змодельовано атаки та запропоновано методи підвищення стійкості до конкретних атак. Оцінено вплив цих атак на стабільність роботи стеганоканалу.

Практична значущість результатів дослідження полягає у дослідженні методу оцінювання стійкості стеганоалгоритмів до поширених атак, що дозволяє здійснювати вибір методу вбудовування на основі характеристик каналу зв'язку та потенційних атак. Як приклад такого вибору було проаналізовано стійкість для методу модифікації кольорового простору та методу модифікації найменш значущих біт, а також при яких атаках варто використовувати той чи інший алгоритм.

Висновки і рекомендації

Досліджено набір методів порівняння стійкості стеганоалгоритмів до атак. Цей метод дозволяє розширити аналіз алгоритмів вбудовування та до наявних характеристик захисту від виявлення, співвідношення сигналу до шуму та структурної подібності додати інваріантність до атак. Це дозволяє виконувати вибір алгоритму більш точно, залежно від характеристик каналу зв'язку та очікуваних атак.

Проведено моделювання розповсюджених у каналах передачі даних геометричних атак, що включають афінні перетворення, стиснення та редагування. Оцінено стійкість стеганографічного повідомлення, яке приховане в контейнері методом модифікації найменш значущих біт (LSB) та методом модифікації кольорового простору (Color space cube modification).

Проаналізовано різницю між вирізанням ряду пікселів та масштабуванням на один піксель, останній перераховуючи значення пікселів унеможливує видобування. Встановлено, що метод LSB стійкіший за метод модифікації кольорового простору у випадку масштабування на 9.26%, адже ця атака вносить суттєві зміни до кольорового простору. Результат оцінки стійкості методу LSB становить 10% та показує, що стійкість проти атаки масштабуванням може бути покращена.

Відзначено перевагу методу кольорового простору, як такого, що не залежить від зміни порядку пікселів, у простих випадках редагування. Рівень цієї переваги коливається в межах 14-98%, залежно від того, наскільки атака впливає на послідовність біт у повідомленні, вбудованому методом LSB.

Виявлено неефективність обох методів до атаки стисненням, оскільки перерахунок значення відбувається чи не для всіх пікселів і вбудовані дані майже повністю втрачені. Запропоновано методи покращення стійкості алгоритмів, а також можливості відновлення повідомлень після таких атак за допомогою розбивання на блоки та кодування з виявленням та виправленням помилок.

Перелік посилань

1. Kalenyuk, P., Rybyska, O., & Ivasyk, G. (2019). Linear algebra and analytic geometry: Basic course: Tutorial [Лінійна алгебра та аналітична геометрія. Базовий курс] (J. Wojtowicz, Trans.). Lviv Polytechnic Publishing House.
2. Walia, Ekta & Jain, Payal & Navdeep. (2010). An analysis of LSB & DCT based steganography. Global Journal of Computer Science and Technology. 10.
3. Zhang, Y., Luo, X., Wang, J., Yang, C., & Liu, F. (2018). A robust image steganography method resistant to scaling and detection. Journal of Internet Technology, 19(2), 607–618. <https://doi.org/10.3966/1607926420180319-02029>.
4. Apau, R., Asante, M., Twum, F., Ben Hayfron-Acquah, J., & Peasah, K. O. (2024). Image steganography techniques for resisting statistical steganalysis attacks: A systematic literature review. PloS one, 19(9), e0308807. <https://doi.org/10.1371/journal.pone.0308807>.
5. Alrusaini, O. A. (2025). Deep learning for steganalysis: Evaluating model robustness against image transformations. Frontiers in Artificial Intelligence, 8, 1532895. <https://doi.org/10.3389/frai.2025.1532895>.

6. AbdelRaouf, A. (2021). A new data hiding approach for image steganography based on visual color sensitivity. *Multimedia Tools and Applications*, 80. <https://doi.org/10.1007/s11042-020-10224-w>.
7. Margalikas, E., & Ramanauskaitė, S. (2019). Image steganography based on color palette transformation in color space. *EURASIP Journal on Image and Video Processing*, 2019(1). <https://doi.org/10.1186/s13640-019-0484-x>.
8. Гасілін Д., Журавель І. (2024) Стеганографічний метод приховування інформації через модифікацію колірного простору з урахуванням властивостей зорового сприйняття. *Інформаційні технології і автоматизація – 2024 : матеріали XVII Міжнародної науково-практичної конференції, Одеса, 31 жовтня – 1 листопада 2024 р. – 2024. – С. 175–178.*
9. Virtanen, P., Gommers, R., Oliphant, T. E., Haberland, M., Reddy, T., Cournapeau, D., Burovski, E., Peterson, P., Weckesser, W., Bright, J., van der Walt, S. J., Brett, M., Wilson, J., Millman, K. J., Mayorov, N., Nelson, A. R. J., Jones, E., Kern, R., Larson, E., Carey, C. J., Polat, İ., Feng, Y., Moore, E. W., VanderPlas, J., Laxalde, D., Perktold, J., Cimrman, R., Henriksen, I., Quintero, E. A., Harris, C. R., Archibald, A. M., Ribeiro, A. H., Pedregosa, F., van Mulbregt, P., & SciPy 1.0 Contributors. (2020). *SciPy 1.0: Fundamental algorithms for scientific computing in Python*. *Nature Methods*, 17(3), 261–272. <https://doi.org/10.1038/s41592-019-0686-2>.
10. Harris, C. R., Millman, K. J., van der Walt, S. J., et al. (2020). Array programming with NumPy. *Nature*, 585, 357–362. <https://doi.org/10.1038/s41586-020-2649-2>.
11. The Pandas Development Team. (2024). *pandas-dev/pandas: Pandas (v2.2.3)*. Zenodo. <https://doi.org/10.5281/zenodo.13819579>.
12. Bradski, G. (2000). The OpenCV library. *Dr. Dobb's Journal of Software Tools*.
13. Clark, A. (2015). Pillow (PIL fork) documentation. Read the Docs. Retrieved from <https://pillow.readthedocs.io/>.
14. Malluri, S. (2021). Image Steganography. GitHub. <https://github.com/LudicrousWhale/ImageSteganography>.
15. Shah, M., Yu, X., Di, S., Becchi, M., & Cappello, F. (2024). A portable, fast, DCT-based compressor for AI accelerators. In *Proceedings of the 33rd International Symposium on High-Performance Parallel and Distributed Computing* (pp. 109–121). Association for Computing Machinery. <https://doi.org/10.1145/3625549.3658662>.
16. Kunthoth, Jayakanth & Subramanian, Nandhini & Al-ma'adeed, Somaya & Bouridane, Ahmed. (2023). Video steganography: recent advances and challenges. *Multimedia Tools and Applications*. 82. 1-43. <https://doi.org/10.1007/s11042-023-14844-w>.
17. Brakensiek, J., Gopi, S., & Makam, V. (2023). Generic Reed-Solomon codes achieve list-decoding capacity. In *Proceedings of the 55th Annual ACM Symposium on Theory of Computing* (pp. 1488–1501). Association for Computing Machinery. <https://doi.org/10.1145/3564246.3585128>.
18. McKiernan, S. (2023). Foundational Techniques for Wireless Communications: Channel Coding, Modulation, and Equalization. *ArXiv*. <https://arxiv.org/abs/2310.13209>.
19. Yergeau, F. (2003, November). UTF-8, a transformation format of ISO 10646. <https://www.rfc-editor.org/info/rfc3629>.

Надійшла 20.04.2025