

**Нужний Сергій Миколайович**

кандидат технічних наук, доцент,

завідувач кафедри комп'ютерних технологій та інформаційної безпеки

Національний університет кораблебудування імені адмірала Макарова, м. Миколаїв, Україна

ORCID: 0000-0002-7706-0453

s.nuzhniy@gmail.com

НЕЙРОМЕРЕЖЕВА МОДЕЛЬ ОЦІНЮВАННЯ РІВНЯ ЗАХИЩЕНОСТІ СКЛАДНОЗАШУМЛЕНОЇ МОВНОЇ ІНФОРМАЦІЇ НА ОСНОВІ СТРУКТУРНОЇ СХЕМИ RII

Анотація. У роботі розглянуто проблему визначення рівня захищеності мовної інформації в умовах дії складних акустичних та віброакустичних завад, коли традиційні показники якості мовлення (SNR, STI, SII, PESQ, STOI) не відображають реальної здатності сучасних алгоритмів реконструкції (HMM, DNN, BayesianNN, GAN) відновлювати семантичний зміст перехоплених сигналів. За низького відношення сигнал/шум і навіть після застосування активних віброакустичних завад значна частка мовних структур залишається доступною для реконструкції, що створює ризик витоку інформації. Відсутність кількісного критерію, який би узгоджував фізичні, лінгвістичні та нейромережеві параметри сигналу та дозволяв оцінювати фактичну можливість семантичного відновлення, визначає ключову науково-технічну проблему роботи. Для її дослідження запропоновано інтегральну модель оцінювання залишкової мовної інформації складнозашумленої мови з урахуванням спектральної структури, контекстних лінгвістичних залежностей і можливостей сучасних реконструкційних систем. Багаторівнева архітектура моделі включає аналітичний спектральний опис сигналу (A_i , μ_i , σ_i^* , $Z_0(f)$, $s(f)$), баєсівсько-марківську трифонну структуру та багатопарову емісійну нейромережу, побудовану на основі CNN, MLP і BayesianNN. Спектральний рівень забезпечує формальний опис енергетичних максимумів і адаптивного згладження; лінгвістичний – відображає ймовірнісні закономірності переходів між трифонами; нейромережевий – інтегрує всі типи ознак і моделює невизначеність емісійних ймовірностей. На основі синтезу цих рівнів сформовано критерій залишкової розбірливості (RII), що кількісно характеризує здатність потенційного перехоплювача відновити зміст повідомлення після дії завад і фільтрації. Визначено порогове значення RII^* , яке інтерпретується як умовна межа між інформаційно небезпечним та інформаційно недостатнім для реконструкції режимами. Запропонована модель може бути використана в системах технічного захисту інформації, випробувальних лабораторіях і комплексах оцінювання ефективності активних віброакустичних завад. Отримані результати формують науково-технічні засади створення інструментальних методів визначення рівня залишкової інформативності мови та підвищення її захищеності.

Ключові слова: складнозашумлений мовний сигнал; індекс залишкової розбірливості; трифонна модель; емісійна нейромережа; BayesianNN.

ВСТУП

Акустичні сигнали, що формуються в реальних умовах, зазнають впливу комплексних завад, які порушують спектральну структуру мовлення та знижують його розбірливість. Обробка складнозашумлених мовних сигналів (під поняттям «складнозашумленої мовної інформації» в роботі розуміється мовний сигнал, у якому відбувається одночасне накладання кількох різномірних типів завад (або завад, створених системами постановки активних завад), що порушують як фізичну, так і



лінгвістичну структуру мовлення та значно ускладнюють його реконструкцію традиційними методами.) є критично важливою для технічного захисту інформації, оскільки сучасні розвідувальні системи застосовують HMM-(Hidden Markov Model (англ.), Прихована Марківська модель), DNN- (Deep Neural Network (англ.), Глибока нейронна мережа) та баєсівські моделі реконструкції [1] – [4], здатні відновлювати зміст повідомлення навіть за умов часткової втрати акустичних компонентів [5] – [7]. Тому оцінювання інформативності мовного сигналу після процедур фільтрації має визначальне значення для встановлення фактичного рівня його вразливості.

Традиційні показники якості мовлення – SNR (Signal-to-Noise Ratio (англ.), відношення сигнал/шум), STI (Speech Transmission Index, Індекс передачі мовлення), SII (Speech Intelligibility Index (англ.), Індекс розбірливості мовлення), PESQ (Perceptual Evaluation of Speech Quality (англ.), Перцептивна оцінка якості мовлення), STOI (Short-Time Objective Intelligibility (англ.), Короточасний об'єктивний індекс розбірливості) – характеризують фізичні або психоакустичні властивості сигналу [8], [9], але не дають змоги оцінити ймовірність відновлення лінгвістичної структури сучасними алгоритмами розпізнавання [1], [2], [10]. У сфері технічного захисту інформації відсутні моделі, здатні кількісно оцінити залишкову інформативність складнозашумленої мови після її очистки, зокрема в частині визначення ризику семантичної реконструкції перехоплених сигналів противником.

Додаткової складності набуває застосування систем постановки активних віброакустичних завад [11], [12], які створюють маскуючі поля різної спектральної структури. Попри їх ефективність як засобу зниження інформативності, інструментальний контроль їх роботи має суттєві обмеження [13] – [15]. Існуючі методи контролю базуються на вимірюванні енергетичних характеристик шуму або загального рівня акустичного впливу, однак не враховують того, що сучасні системи розпізнавання можуть частково компенсувати дію таких завад та відновити трифонну структуру мовлення [4], [16] – [18]. У результаті неможливо об'єктивно визначити, чи забезпечують активні віброакустичні завади реальну конфіденційність мовної інформації, оскільки ключовим параметром є не рівень шуму як такий, а залишкова здатність сигналу до смислової реконструкції.

Таким чином, виникає необхідність створення інтегральної моделі оцінювання рівня захищеності мовної інформації, яка враховує характер завад, параметри процесів фільтрації, спектральні та лінгвістичні властивості мовлення, а також ймовірнісні механізми реконструкції, притаманні сучасним нейромережевим системам. Запропонований підхід реалізується через показник реконструктивної інформативності RII (Residual Intelligibility Index (англ.), індекс залишкової розбірливості), що пов'язує аналітичні спектральні параметри, марківсько-баєсівську трифонну модель та нейромережеву емісійну підсистему і дозволяє кількісно оцінити ймовірність відновлення змісту сигналу.

Метою дослідження є розроблення науково обґрунтованої моделі оцінювання рівня захищеності мовної інформації на основі визначення залишкової реконструктивної інформативності складнозашумленої мови після її фільтрації та дії активних віброакустичних завад.

Для досягнення поставленої мети необхідно розв'язати такі основні задачі:

1. Синтезувати інтегровану багаторівневу математичну модель визначення реконструктивної інформативності складнозашумленого мовного сигналу (RII), що включає аналітичну модель спектральних параметрів, баєсівсько-марківську трифонну структуру та нейромережеву реконструкційну підсистему.



2. Розробити методику інструментального контролю рівня захищеності мовної інформації та ефективності систем активних віброакустичних завад на основі інтегрального критерію RII.

АНАЛІЗ ОСТАННІХ ДОСЛІДЖЕНЬ І ПУБЛІКАЦІЙ ЗА ТЕМОЮ

За останні 5–7 років помітно зросла увага до задачі обробки та реконструкції мовлення в умовах низького відношення сигнал/шум (low SNR) [14], [15], [19], насамперед у контексті телекомунікацій, систем розпізнавання мовлення та безпекових застосувань. У сучасних оглядових роботах відзначається перехід від класичних методів спектрального віднімання, Wiener-фільтрації та статистичних моделей до глибоких нейромережових архітектур [1], [20], [21], включно з автоенкодерами, CNN (Convolutional Neural Network (англ.), Згортоква нейронна мережа), RNN (Recurrent Neural Network (англ.), Рекурентна нейронна мережа), трансформерами та гібридними моделями, що одночасно орієнтовані на покращення якості сигналу й підвищення розбірливості для людини та ASR-систем (Automatic Speech Recognition (англ.), Автоматичне розпізнавання мовлення).

У низці досліджень останніх років фокусується увага саме на наднизьких значеннях SNR, де класичні алгоритми практично втрачають ефективність. Так, у роботах зі створення спеціалізованих глибоких мереж для low-SNR середовищ пропонуються комплексні трансформерні архітектури в комплексній площині (complex domain) [18], [19], здатні моделювати довготривалі часові залежності та підвищувати як об'єктивні метрики (PESQ, STOI), так і суб'єктивну розбірливість мовлення при дуже низькому SNR.

Окремі роботи спрямовані на побудову двоетапних мереж із модулями шумопригнічення та відновлення проміжно «перерізаних» компонент мовного сигналу, орієнтованих саме на екстремальні умови завад [14].

Додатково розробляються моделі для низько-SNR середовищ із використанням детальних спікер-ембедингів (компактне числове представлення (вектор), яке описує унікальні голосові характеристики конкретної людини), які демонструють помітне зростання точності ASR у складних умовах [22].

Суттєвий блок англійських робіт присвячений інтеграції задачі підсилення/очищення мовлення й задачі розпізнавання [10], [17], [23]. У систематичних оглядах deep-learning-методів (методів глибокого навчання) підкреслюється, що більшість сучасних архітектур (DNN, CNN, LSTM (Long Short-Term Memory (англ.), Довготривала короткочасна пам'ять), трансформери) одночасно оптимізуються за критеріями якості мовлення (PESQ, STOI) та показниками помилки розпізнавання (WER), що фактично реалізує оцінювання «реконструктивної здатності» алгоритмів.

Окремі дослідження показують, що в діапазоні SNR 0...5 дБ людина втрачає сприйняття змісту [8], [9], тоді як глибокі моделі все ще здатні класифікувати мовні й небезпечні акустичні події з прийнятною точністю, що суттєво змінює уявлення про «фактичну» залишкову інформативність сигналу в безпекових системах.

Паралельно активно розвиваються генеративні та байєсівські підходи до реконструкції мовлення [4], [5], [7], [24]. Роботи з використанням GAN-архітектур (Generative Adversarial Network (англ.), Генеративна змагальна мережа) для перетворення шепітного мовлення у «озвучене» демонструють можливість відновлення голосових компонент, які несуть суттєву частку семантичної інформації, навіть за відсутності повноцінного спектра основного тону.



Додатково розробляються аудіо-візуальні моделі, де до акустичних ознак додаються відеодані (рух губ), що істотно підвищує якість реконструкції при дуже низькому SNR; такі системи прямо орієнтовані на сценарії спостереження та безпеки [25]

В оглядах з deep learning-based speech enhancement (покращення мовлення на основі глибокого навчання) та Bayesian deep learning (Баєсівське глибоке навчання) окреслюється тренд на використання варіаційних, ймовірнісних моделей, що явно описують невизначеність оцінок та стають базою для баєсівських реконструкційних схем у мовних технологіях [3], [4].

Окремий напрямок становлять дослідження мовної приватності та акустичного маскування в офісних просторах і навчальних середовищах. Міжнародні стандарти ISO 3382-3:2022 та ISO 22955:2021, а також прикладні роботи з акустики open-plan офісів описують підходи до кількісної оцінки рівня мовної приватності через STI, просторовий спад мовлення D_2, S , відстані відволікання та приватності r_D , r_P , а також через параметри фонових завад і звукомаскуючих систем [11], [12], [26].

У низці досліджень і професійних гайдів наголошується на ролі електронного звукомаскування, правильно налаштованого спектра маскуючих шумів та конструктивних параметрів офісних перегородок для досягнення прийняттого рівня конфіденційності мовлення; при цьому ключовими залишаються індекси STI/SII, а не безпосередня ймовірність реконструкції мовлення машинними системами.

Незважаючи на суттєвий прогрес у сфері підсилення мовлення, систем покращення мовлення з низьким співвідношенням сигнал/шум, генеративних моделей та акустичної приватності, аналіз англійських джерел показує, що переважна більшість робіт орієнтована на покращення якості та розбірливості мовного сигналу, а не на оцінювання рівня його захищеності. Сучасні публікації фокусуються на мінімізації WER, підвищенні PESQ/STOI, покращенні класифікації в умовах шуму або забезпеченні нормативних значень STI та акустичних параметрів у приміщеннях. Водночас не пропонується інтегральної моделі, яка б прямо зв'язувала параметри низько-SNR мовного сигналу, властивості реконструкційних алгоритмів (HMM/DNN/Bayesian/GAN) та показники просторово-акустичного середовища в єдиний критерій “залишкової реконструктивної інформативності”, придатний для кількісного визначення рівня захищеності мовної інформації. Саме ця прогалина у сучасних англійських дослідженнях і обумовлює необхідність розроблення моделі типу RII, орієнтованої не лише на якість відновлення, а й на оцінювання ймовірності семантичної реконструкції мовлення в контексті технічного захисту інформації.

ЗАГАЛЬНА ІДЕОЛОГІЯ МОДЕЛІ

Ідеологія баєсівсько-марківської нейромережевої моделі оцінювання рівня захищеності мовної інформації базується на об'єднанні трьох фундаментальних рівнів опису мовного сигналу: фізичного (акустичного), мовного (лінгвістичного) та інформаційного (канального). На відміну від класичних моделей оцінювання розбірливості, які використовують суто енергетичні метрики на кшталт SNR, STI, SII або SPC, запропонована архітектура розглядає мовлення як структуру смислових одиниць, здатних відтворюватися засобами технічної розвідки навіть при сильних завадах. Це означає, що критерієм стає не чутність, а ймовірність коректної реконструкції цілісного повідомлення [2], [10].

У центрі моделі знаходиться трифон – контекстно залежна одиниця мовлення (X , Z , Y), що враховує вплив сусідніх фонем. Саме трифонна структура визначає семантичну

стійкість мовлення до завад, тому модель оцінює імовірність відновлення не окремих звуків, а їх семантичної послідовності. Такий підхід дозволяє зіставити фізичні властивості сигналу зі структурними закономірностями мови. На рис.1 приведена функціональна схема оцінювання рівня залишкової розбірливості слів

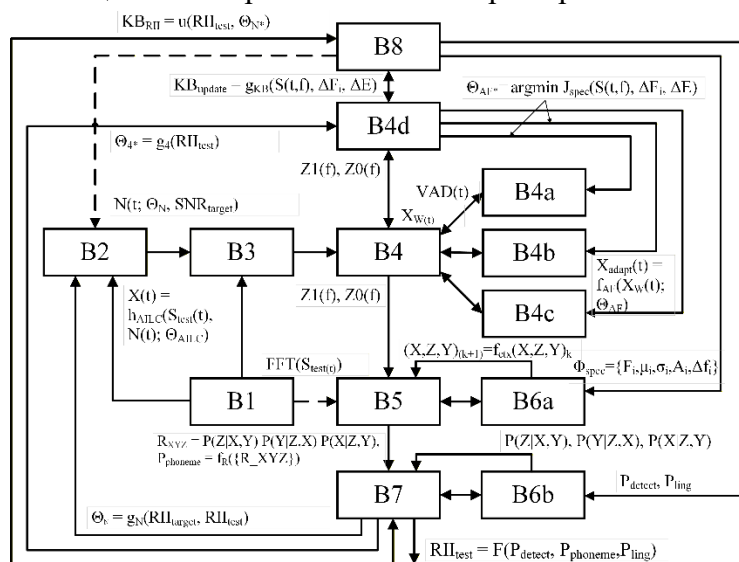


Рис. 1. Функціональна схема оцінювання рівня захищеності складнозашумленої мовної інформації на основі структурної схеми RII

На рис.1 прийняті позначення:

B1 – блок формування тест-сигналу; B2 – блок формування сигналу завади; B3 – блок формування сигналу в технічному каналі витоку інформації (AIRC); B4 – блок попередньої обробки; B4a – виявлення пауз; B4b – вейвлет-фільтрація; B4c – адаптивна фільтрація; B4d – спектральний аналіз та оптимізація коефіцієнтів; B5 – спектральний + трифонний аналіз; B6a – аналіз трифонів попередній; B6b – аналіз трифонів основний; B7 – аналіз залишкової розбірливості; B8 – база знань.

Далі вводиться баєсівська концепція ймовірнісної оцінки параметрів. Це дає змогу моделювати невизначеність як самого мовного сигналу, так і каналу. Будь-який технічний канал має власні характеристики (реверберацію, амплітудно-частотну нерівномірність, власний шум, спектр завад), які не можуть бути представлені детерміновано.

Баєсівський підхід дозволяє моделювати розподіли параметрів, а не окремі величини, що значно підвищує стійкість системи.

Ключовий елемент – ієрархічність. На першому рівні формується спектральний опис мовлення: піки, форманти, енергетичні центри, ширини смуг, логарифмічні огибаючі. На другому рівні відтворюється лінгвістична структура шляхом побудови марківських залежностей між трифонами. На третьому рівні оцінюється вплив каналу на можливість відтворення смислової послідовності. На четвертому – формується інтегральний індекс реконструкції RII.

Узагальнюючи, ідеологія моделі формує ймовірнісний опис процесу: «від мовного джерела → до спотвореного сигналу → до оцінки можливості смислового відновлення».

Це дозволяє прогнозувати здатність технічних засобів розвідки до дешифрування мовних повідомлень навіть у складних акустичних і технічних умовах.

МАРКІВСЬКА МОДЕЛЬ (ПММ)

Прихована марківська модель (ПММ) формує лінгвістичний рівень опису системи (див. рис.2). Для української мови характерні сильні контекстні зв'язки, тому використання трифонів забезпечує природний механізм прогнозування мовних структур. Станами ПММ є трифони $T_k = (X_k, Z_k, Y_k)$, а переходи визначаються ймовірностями: $P(T_k | T_{k-1})$, які отримуються з корпусної статистики та згладжуються баєсівським чином через апіорі Дирихле, що запобігає появі нульових частот.

Аналітичні залежності, взяті з [27], виконують роль кількісних мір:

$$R_{XYZ} = P(Z|X, Y) \cdot P(Y|Z, X) \cdot P(X|Z, Y)$$

де $P(Z|X, Y)$, $P(Y|Z, X)$, $P(X|Z, Y)$ – складові, що відображають ступінь узгодженості трифона з його контекстом.

У разі деградованих сигналів саме висока узгодженість дає змогу припускати реальну мовну структуру навіть при неповних або спотворених акустичних даних.

ПММ у моделі виконує роль «мовного стабілізатора»: вона не дозволяє емісійній моделі видавати неможливі послідовності, тим самим різко знижуючи кількість хибних реконструкцій. Навіть коли окремі спектральні ознаки сильно пошкоджені (перекриті шумом, завадами чи реверберацією), ПММ здатна вибрати найбільш правдоподібний шлях у просторі фонетичних послідовностей. Це є критично важливим для оцінки РІІ.

Особливість української мови – велика кількість коартикуляційних зсувів. Вони враховуються через матриці переходів, побудовані на реальному корпусі мовлення. Таким чином, модель отримує статистичні правила переходу між станами, які відповідають природному мовленню і не допускають штучних комбінацій [2].

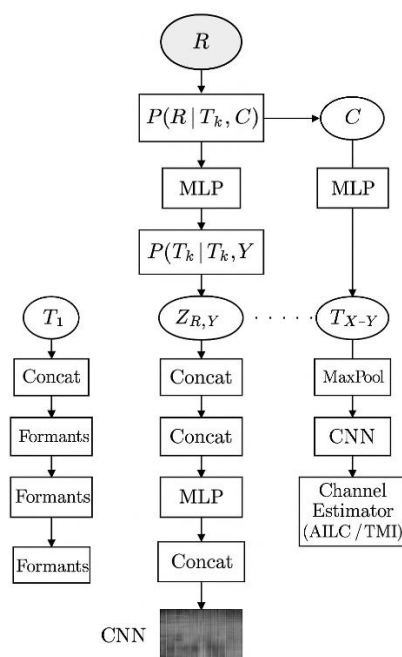


Рис. 2. Баєсівсько-марківська архітектура моделі.



На рис. 2 прийняті позначення:

R – випадкова змінна, що задає апостеріорну ймовірність реконструкції (Residual Informativeness Prior);

P(R | T_k, C) – блок обчислення умовної ймовірності відновлення R за контекстом трифона T_k та класом C;

C – класовий контекст (Channel/Condition), що описує акустичні умови та тип завади;

MLP – багатошаровий перцептрон (Multilayer Perceptron), використаний для параметричного моделювання ймовірностей та емісійних векторів;

P(T_k | T_{k-1}, Y) – блок обчислення ймовірності трифонного переходу, що моделює марковську залежність між трифонами (Triphone Transition Model);

Z_{R,Y} – латентний вектор (емісійний embedding), що поєднує апостеріорні R та контекстні ознаки Y;

T₁ – перший трифон або лінгвістичний контекст початкової позиції (Triphone Start-State);

Concat – блок конкатенації ознак, що об'єднує спектральні, лінгвістичні та нейромережіві представлення;

Formants – блок оцінювання формантної структури сигналу (F₁–F₄ Estimator);

T_{x-y} – діапазон трифонів, що відповідає локальному контекстному вікну;

MaxPool (MP) – блок субсемплювання на основі операції максимізації (Max Pooling);

CNN – згортова нейронна мережа (Convolutional Neural Network), що моделює локальні спектральні патерни;

Channel Estimator (AILC / TMI) – блок класифікації або регресії каналу (AILC – Attacker Information Leakage Capacity; TMI – Total Message Intelligibility), визначає рівень інформаційної доступності мовного сигналу після реконструкції.

БАЄСІВСЬКА ЕМІСІЙНА МОДЕЛЬ

Емісійна модель визначає залежність між кутовими спостережуваними ознаками F_k і прихованими трифонами T_k . Вона складається з двох гілок нейромережі: MLP для формальних спектральних ознак і CNN для спектрограм. Усі параметри моделі тренуються як випадкові змінні за допомогою варіаційного виведення, що дозволяє описувати невизначеність.

Аналітичні залежності, з [27], включено у внутрішні параметри моделі:

1) Поріг Донного:

$$T = \kappa \cdot \sigma_n \cdot \sqrt{2 \ln N} \quad (1)$$

де: κ – коефіцієнт масштабу порогу Донного; σ_n – середнє квадратичне відхилення шуму; N – розмір спектральної вибірки.

2) Критерій злиття піків:

$$|f_i - f_j| < 0.075 \cdot \min(f_i, f_j) \quad (2)$$



де: f_i, f_j – частоти піків спектра; $\min(f_i, f_j)$ – мінімальна з двох частот, 0.075 – емпірично встановлений поріг відносної близькості частот, за якого два піки розглядаються як один фізичний максимум.

3) Енергетичний центр:

$$\mu = \frac{\sum f_k A_k}{\sum A_k} \quad (3)$$

де: f_k – частота та амплітуда k -того спектрального компонента, відповідно.

4) Спектральна модель:

$$S(f) = \sum A_i \exp\left(-\frac{(f - F_i)^2}{2\sigma_i^2}\right) + \lambda \cdot \mu \quad (4)$$

де: A_i – амплітуда i -того піку; F_i – центральна частота піку; σ_i – ширина спектральної смуги; λ – коефіцієнт впливу енергетичного центру μ .

MLP отримує набір ознак: пікову структуру, формантні відхилення, енергетичні центри, відношення сигнал/шум, а CNN – двовимірну матрицю спектрограми. Після обробки вони зливаються у латентний простір h_k , який подається у баєсівський блок вибору апостеріорного розподілу.

Таким чином емісійна модель формує не одномірні оцінки, а ймовірнісний опис відповідності «акустика → зміст». Це дозволяє системі бути стійкою до сильних спотворень сигналу і правильно співвідносити неповні або частково знищені ознаки з лінгвістичною структурою трифонів.

АНАЛІТИЧНА МОДЕЛЬ СПЕКТРАЛЬНИХ ПАРАМЕТРІВ

Аналітична модель спектральних параметрів описує фізичну структуру мовного сигналу в частотній області та задає форму входних ознак для баєсівсько-нейромережевої підсистеми [27], [28], [29]. Вихідним об'єктом аналізу є дискретний часовий сигнал $x(t)$, для якого обчислюється модуль спектра за перетворенням Фур'є

$$Z_1(f) = |F\{x(t)\}| \quad (5)$$

де: $x(t)$ – реальний мовний сигнал у часовій області; $F\{\cdot\}$ – оператор дискретного перетворення Фур'є; $Z_1(f)$ – спектральна амплітуда до процедур фільтрації та згладження.

Далі застосовується вейвлет-фільтрація з порогом Донного (див (1)) та, після фільтрації, виконується детекція локальних максимумів $Z_1(f)$ і їх агрегація за критерієм (2).

На основі знайдених максимумів формується гаусівська огинаюча



$$Z_0(f) = \sum A_i \cdot \exp\left(-\frac{(f - \mu_i)^2}{2\sigma_i^2}\right) \quad (6)$$

де: A_i – амплітуда i -го гаусівського компонента, що відображає внесок відповідного максимуму; μ_i – частотний центр компонента, пов'язаний із положенням локального максимуму; σ_i – ширина компонента, яка інтерпретується як ефективна смуга концентрації енергії навколо μ_i ; Σ – сумування ведеться за всіма виділеними максимумами спектра.

Для стабілізації поведінки огинаючої у високочастотній області використовується логарифмічна згладжувальна функція

$$s(f) = a \cdot \log(1 + b \cdot f) \quad (7)$$

де: a – параметр масштабу згладженості, що задає загальний рівень пом'якшення нерівностей спектра; b – параметр, який визначає швидкість зростання згладжувального ефекту з частотою; f – поточне значення частоти у досліджуваному діапазоні.

Корекція ширин компонент виконується через глобальний енергетичний центр спектра

$$\sigma_i^* = \sigma_i \cdot f(\mu) \quad (8)$$

$$\mu = \frac{\sum (f_k \cdot Z(f_k))}{\sum (Z(f_k))} \quad (9)$$

де: μ – центр ваги спектральної енергії, що показує, в якій області частот зосереджена основна енергія сигналу; f_k – дискретні частотні відліки; $Z(f_k)$ – значення спектра на відповідних частотах; $f(\mu)$ – безрозмірна функція, яка модифікує σ_i залежно від положення μ , забезпечуючи узгодженість між локальними і глобальними характеристиками спектра.

У сукупності залежності $Z(f)$, λ , критерій агрегації піків, гаусівська модель $Z_0(f)$, функція $s(f)$ та корекція σ_i^* утворюють повний спектральний апарат моделі, що використовується для формування вектора ознак F_k на вході нейромережевої підсистеми реконструкції. Це співвідношення використовується для побудови числового алгоритму оцінювання параметрів моделі у задачах реконструкції фрикативних фонем. Такий запис дозволяє виокремити незалежні змінні та параметри, що полегшує процедуру ідентифікації моделі за експериментальними даними. Формула безпосередньо використовується на етапі навчання нейромережевої частини, де параметри вбудовуються у вектор ознак F_k .

Використання саме цієї форми рівняння спрощує аналітичний вивід та дозволяє отримати замкнуті вирази для похідних при оптимізації. Параметри, що входять у дану залежність, оцінюються стійкими методами, що мінімізують вплив аномальних значень та шумових викидів. На практиці це дає змогу реалізувати обчислювальні процедури в реальному часі без істотного зростання складності алгоритму.

Таке формулювання також є зручним для подальшого розширення моделі за рахунок введення додаткових параметрів або коригувальних коефіцієнтів.

БАЄСІВСЬКО-МАРКІВСЬКА ТРИФОННА МОДЕЛЬ

Баєсівсько-марківська трифонна модель (див. рис.3) описує зміну фонематичних станів у часі з урахуванням контексту та спектральних ознак. Основною одиницею є трифон

$$T_k = (X_k, Z_k, Y_k) \quad (10)$$

де: X_k – попередня фонема; Z_k – центральна фонема; Y_k – наступна фонема; k – індекс позиції у мовній послідовності.

Перехід між станами задається марківською залежністю першого порядку

$$P(T_k | T_{\{k-1\}}) = P(X_k, Y_k | Z_{\{k-1\}}) \quad (11)$$

де: $P(T_k | T_{\{k-1\}})$ – ймовірність переходу від попереднього стану до поточного; $P(X_k, Y_k | Z_{\{k-1\}})$ – умовна ймовірність появи контексту (X_k, Y_k) для заданої попередньої центральної фонем $Z_{\{k-1\}}$.

Емісійна ймовірність описується залежністю

$$P(F_k | T_k, C) = \text{BayesianNN}(F_k, e_C), \quad (12)$$

де: F_k – вектор спектральних ознак, сформований на основі параметрів $A_i, \mu_i, \sigma_i^*, s(f)$; C – каналний стан, що характеризує умови поширення сигналу; e_C – embedding-вектор каналу (Ембедінг - це числове (векторне) подання семантики (сенсу) тексту), який кодує SNR, SPC, SII та інші параметри середовища; $\text{BayesianNN}(\cdot)$ – баєсівська нейромережа, яка повертає умовний розподіл ймовірностей F_k для кожного трифона T_k .

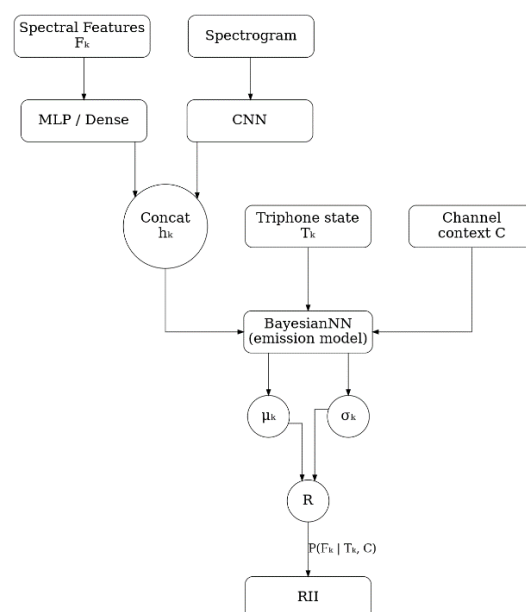


Рис. 3. Баєсівсько-марківська нейромережева модель спектрально-трифонної реконструкції індексу інформативності RII



На рис. 3 прийняті позначення:

Spectral Features F_k — вектор спектральних ознак для k -го фрейма, отриманий із нормованого спектра або огибаючої спектра.

MLP / Dense — повнозв'язний шар або багат шарова перцептронна компонента, що виконує нелінійне перетворення спектральних ознак F_k .

Spectrogram — матриця спектрограми (STFT), що використовується як вхід до згорткової гілки.

CNN — згорткова нейромережа для виділення локальних часово-частотних ознак зі спектрограми.

Concat h_k — оператор конкатенації; формує латентний вектор h_k з об'єднання виходів спектральної (MLP) і згорткової (CNN) гілок.

Triphone state T_k — трифонний стан k -го такту марковського ланцюга (HMM), що відображає контекстну залежність між фонемами.

Channel context C — каналові параметри умов передачі/фільтрації, що впливають на реконструкцію (тип завади, SNR, спектральна характеристика тощо).

BayesianNN (emission model) — баєсівська нейромережа, що моделює емісійний розподіл ознак та повертає параметри ймовірнісної моделі.

μ_k — математичне сподівання емісійного розподілу для k -го фрейма.

σ_k — дисперсійний/коваріаційний параметр емісійного розподілу для k -го фрейма.

R — вузол формування локальної оцінки залишкової інформативності на основі (μ_k, σ_k) .

RII — інтегральна метрика залишкової реконструктивної інформативності, що оцінює ймовірну здатність відновлення мовного фрагмента при відомих T_k та C .

Апостеріорний розподіл станів визначається за forward–backward алгоритмом

$$P(T_k | F_{1:K}, C) = \alpha_k \cdot \beta_k \quad (13)$$

де: $F_{1:K}$ — послідовність векторів ознак для всіх часових індексів; α_k — коефіцієнти прямого проходу (forward), що акумулюють ймовірності від початку послідовності до k ; β_k — коефіцієнти зворотного проходу (backward), що враховують внесок наступних спостережень.

Статистика трифонів за корпусом описується виразом

$$P(X, Y | Z) = \frac{N_{\{XYZ\}}}{N_Z} \quad (14)$$

де: N_Z — загальна кількість появ центральної фонемі Z у корпусі; $N_{\{XYZ\}}$ — кількість появ трійки (X, Z, Y) ; $P(X, Y | Z)$ — відносна частота реалізації контексту (X, Y) для заданої центральної фонемі Z .

Акустико-лінгвістичний індекс локальної узгодженості задається як

$$R_{\{XYZ\}} = P(Z | X, Y) \cdot P(Y | Z, X) \cdot P(X | Z, Y), \quad (15)$$

де: $P(Z | X, Y)$ — ймовірність появи центральної фонемі Z при фіксованому контексті (X, Y) ; $P(Y | Z, X)$ — ймовірність наступної фонемі Y за умови (Z, X) ; $P(X | Z, Y)$ — ймовірність



попередньої фонемі X за умови (Z, Y) ; $R_{\{XYZ\}}$ – числова міра узгодженості трійки, що враховує всі напрямки умовних залежностей.

Таким чином, трифонна модель поєднує статистичні властивості корпусу, каналні умови та спектральні ознаки, створюючи повний ймовірнісний опис послідовності мовних станів, який далі використовується у формуванні індексу РІІ.

НЕЙРОМЕРЕЖЕВА АРХІТЕКТУРА РЕКОНСТРУКЦІЇ РІІ

Нейромережева архітектура інтегрує спектральні, каналні та трифонні ознаки у єдину модель, що прогнозує індекс інформативності РІІ. На вході формується спектральний вектор F_k і embedding каналу e_C . Далі ці величини потрапляють до багаторівневої мережі, що містить згорткові та повнозв'язані шари, а також баєсівські компоненти для моделювання невизначеності.

На рівні інтерпретації результатів важливою є оцінка найімовірнішого стану

$$r_k = \max P(T_k | F_{\{1:K\}}, C) \quad (16)$$

де: r_k – індекс (або вектор), що відповідає найімовірнішому трифонному стану T_k ; $P(T_k | F_{\{1:K\}}, C)$ – апостеріорний розподіл для k -го кроку, обчислений за допомогою forward-backward та емісійної моделі; $\max$ – операція вибору стану з максимальною ймовірністю.

Для навчання всієї системи вводиться функція втрат

$$L = L_{HMM} + \lambda_1 \cdot L_{RII} + \lambda_2 \cdot KL \quad (17)$$

де: L_{HMM} – складова, яка штрафує невідповідність між передбаченими та фактичними послідовностями станів у марківській моделі; L_{RII} – складова, що відображає похибку у прогнозуванні індексу РІІ відносно цільових еталонних значень; KL – термін Kullback–Leibler дивергенції, який реалізує баєсівську регуляризацию параметрів мережі; λ_1, λ_2 – вагові коефіцієнти, що задають відносну важливість складових L_{RII} та KL у загальній функції L .

Фінальна оцінка індексу інформативності визначається як

$$RII_{hat} = MLP(r_1, \dots, r_K, e_C), \quad (18)$$

де: RII_{hat} – модельна (прогнозована) оцінка індексу; $MLP(\cdot)$ – багатошарова перцептронна структура, яка обробляє послідовність r_k разом із embedding e_C ; $r_1, \dots, r_K$ – послідовність найімовірніших станів для всіх часових індексів; e_C – вектор, що кодує характеристику каналу.

Узгоджена робота рівнянь для r_k , L та RII_{hat} забезпечує спільне навчання параметрів, при якому спектральна, лінгвістична і канална інформація враховується одночасно, що є ключовим для побудови стійкого критерію захищеності.

ІНТЕГРАЛЬНИЙ КРИТЕРІЙ ЗАХИЩЕНОСТІ МОВНОЇ ІНФОРМАЦІЇ (РІІ)

Індекс РІІ визначає здатність системи до семантичної реконструкції мовлення в умовах дії завад [29]. Він інтегрує спектральні характеристики сигналу, статистику



трифонів та показники каналу в єдину кількісну метрику. Формально індекс задається залежністю

$$RII = g(SNR, SPC, SII, \sigma_i^*, \mu, R_{XYZ}, e_c) \quad (19)$$

де: $g(\cdot)$ – нелінійна вектор-функція, яка реалізується нейромережевою головою і може враховувати як лінійні, так і високого порядку взаємодії між аргументами; SNR – відношення сигнал/шум, що характеризує ступінь деградації спектра; SPC – просторовий показник конфіденційності, який пов'язаний із геометрією приміщення та розташуванням джерел і приймачів; SII – індекс розбірливості мовлення, що відображає можливість правильно розпізнавати слова за умов дії завад; σ_i^* – адаптивні ширини спектральних компонентів, які несуть інформацію про розмивання формантних зон; μ – глобальний енергетичний центр спектра, чутливий до зміщень частотної області концентрації енергії; $R_{\{XYZ\}}$ – акустико-лінгвістичний показник узгодженості трифонної структури; e_c – embedding каналу, в якому агрегуються додаткові параметри середовища, не представлені явно у SNR , SPC і SII .

Практичний сенс критерію RII полягає в тому, що значення індексу безпосередньо інтерпретується як міра інформативності мовного сигналу для потенційного перехоплювача. Порогове значення RII^* задає межу між конфіденційним та неконфіденційним режимами роботи каналу:

- якщо $RII < RII^*$, мовний сигнал вважається нереконструйованим у смисловому відношенні, навіть якщо він частково чутний;
- якщо $RII \geq RII^*$, існує ненульова ймовірність відновлення змісту повідомлень.

Таким чином, введення інтегрального індексу RII дозволяє узагальнити класичні енергетичні та розбірливісні критерії, переформулювавши задачу контролю захищеності як задачу оцінювання ймовірності успішної семантичної реконструкції мовної інформації в умовах дії завад і каналних спотворень.

ВИСНОВКИ

У роботі розв'язано задачу кількісного оцінювання рівня захищеності складнозашумленої мовної інформації шляхом аналізу залишкової реконструктивної інформативності мовного сигналу. Уперше обґрунтовано, синтезовано та реалізовано нейромережеву баєсівсько-марківську модель визначення залишкової інформативності, що інтегрує спектральні, лінгвістичні та каналні параметри в індекс RII . Отримані результати мають науково-методологічне значення, забезпечуючи формальний механізм оцінювання можливості семантичної реконструкції мовлення в умовах комплексних завад.

У межах першої дослідницької задачі описано математичну структуру багаторівневої моделі аналізу складнозашумленого мовлення. Аналітичний рівень включає спектральні залежності A_i , μ_i , σ_i^* , огинаючу $Z_o(f)$, адаптивну функцію згладжування $s(f)$ та механізми ідентифікації енергетичних максимумів. Лінгвістичний рівень реалізовано через баєсівсько-марківську трифонну модель з умовними та спільними ймовірностями переходів між контекстними одиницями. Реконструкційний рівень представлено нейромережею емісійного типу, яка виконує узгодження фізичних і лінгвістичних ознак та забезпечує обчислення RII як кількісної характеристики можливості реконструкції сигналу.



У межах другої дослідницької задачі сформовано науково-методологічний підхід до інструментального контролю рівня захищеності мовної інформації з урахуванням дії активних віброакустичних завад. Основою підходу є багатошарова емісійна нейромережева модель, що включає згорткові шари (2–3 Conv-рівні, 32–128 фільтрів) для виділення стійких спектрально-часових структур, повнозв'язні шари (2–4 MLP-рівні, 128–512 нейронів) для інтеграції акустичних і контекстних параметрів, а також баєсівські шари (1–2 BayesianNN-рівні) для оцінювання невизначеності емісійних ймовірностей. Така архітектура забезпечує адаптацію моделі до різних типів складнозашумлених сигналів і підвищує точність визначення РП за рахунок коректнішої інтерпретації збережених інформативних компонентів мовлення. На основі моделі визначено порогове значення РП*, яке використовується як умовний критерій переходу між режимами інформаційної вразливості та інформаційної недостатності для семантичного відновлення.

Отримані результати мають науково-методологічне значення для побудови систем контролю рівня захищеності мовних каналів. Показник РП та відповідні процедури оцінювання можуть застосовуватися у випробувальних лабораторіях, системах технічного захисту інформації, засобах оцінювання ефективності активних віброакустичних завад і в комплексах аналізу уразливості мовних сигналів. Підхід на основі реконструктивної інформативності створює основу для подальшого розвитку методів формалізації критеріїв мовної безпеки та вдосконалення засобів зниження ризику семантичної реконструкції мовної інформації.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Deng, L., & Li, X. (2020). Deep learning in speech processing: A review. *IEEE Signal Processing Magazine*, 37(3), 107–139.
2. Kolossa, D., & Haeb-Umbach, R. (2018). *Robust Speech Recognition: A Probabilistic Approach*. Academic Press.
3. Pardo, J. M., et al. (2023). Bayesian deep learning for acoustic uncertainty modelling. *IEEE SPL*, 30, 642–646.
4. Zhao, Y., & Ma, J. (2020). Probabilistic reconstruction of masked speech using Bayesian neural networks. *Neural Networks*, 132, 229–241.
5. Arora, A., & Singh, R. (2021). Whisper-to-speech reconstruction using GANs. *Pattern Recognition Letters*, 152, 62–70. <https://doi.org/10.1016/j.patrec.2021.09.011>.
6. Koizumi, Y., et al. (2020). Speech enhancement using deep generative models. *IEEE/ACM TASLP*, 28, 1778–1788.
7. Pascual, S., Ravanelli, M., & Serrà, J. (2020). SEGAN and its descendants: Generative enhancement of degraded speech. *IEEE SPS Magazine*, 37(6), 22–38.
8. Fogerty, D. (2017). Predicting speech intelligibility in noise. *Attention, Perception, & Psychophysics*, 79, 333–344.
9. Wang, D., & Chen, J. (2018). Supervised speech separation based on deep learning. *IEEE TASLP*, 26(10), 1872–1892.
10. Zolfaghari, M., & Sameti, H. (2023). End-to-end noise-robust ASR integrating spectral and linguistic features. *Computer Speech & Language*, 77, 101439.
11. ISO 22955:2021. Acoustics — Acoustic quality of open office spaces.
12. ISO 3382-3:2022. Acoustics — Measurement of room acoustic parameters — Part 3: Open public offices.
13. Reddy, C. K., et al. (2019). The Interspeech Deep Noise Suppression Challenge. *Interspeech*.
14. Li, A., et al. (2021). Two-stage deep enhancement for ultra-low SNR conditions. *Applied Acoustics*, 178, 108018.
15. Bhat, S., & Chatterjee, S. (2020). Complex-domain deep networks for speech enhancement at low SNR. *Speech Communication*, 122, 48–62.



16. Adiga, V. S., & Seltzer, M. L. (2020). End-to-end models for robust speech recognition in low-resource and noisy conditions. *ICASSP*.
17. Dong, L., & Xu, B. (2022). Noise-robust ASR using factorized time-domain convolutional networks. *Speech Communication*, 139, 15–27.
18. Sriram, J., & Chen, M. (2021). Transformer architectures for robust speech recognition. *ICASSP*.
19. Biswas, S., & Manocha, D. (2023). Transformer-based low-SNR speech enhancement using multi-scale spectral attention. *IEEE SPL*, 30, 124–128.
20. Hershey, J. R., et al. (2016). Deep clustering: Discriminative embeddings for segmentation and separation. *ICASSP*.
21. Qin, K., & Wang, D. (2020). Time-domain speech enhancement using deep learning: A survey. *Speech Communication*, 127, 1–16.
22. Andronic, E., & Dehak, N. (2019). End-to-end text-independent speaker verification for extremely noisy environments. *Interspeech 2019*, 4075–4079. <https://doi.org/10.21437/Interspeech.2019-2447>
23. Chen, J., Wang, X., & Xu, Y. (2021). Multi-task neural networks for joint speech enhancement and ASR. *IEEE TASLP*, 29, 2071–2084.
24. Kim, J., & Hahn, M. (2018). Voice conversion using GANs. *NeurIPS Workshop*.
25. Ma, J., & Zhao, Y. (2022). Audio-visual speech reconstruction in extremely noisy environments. *IEEE TPAMI*, 44(8), 4105–4117.
26. Alkin, B. (2021). *Applied acoustics research for speech privacy*. Springer. <https://doi.org/10.1007/978-3-030-67379-8>. ISBN 978-3-030-67378-1
27. Нужний, С. (2025). Удосконалення алгоритму відновлення лінгвістичної складової мовної інформації фонемно-формантним методом для задач оцінки рівня її захищеності. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 3(27). <https://doi.org/10.28925/2663-4023.2025.27>
28. Нужний, С. (2025). Адаптивний фонемно-спектральний метод відновлення мовлення в умовах активних акустичних завад. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 4(28). <https://doi.org/10.28925/2663-4023.2025.28.794805>
29. Нужний, С. (2025). Оцінювання рівня захищеності складнозашумленої мовної інформації за критерієм залишкової розбірливості мови. *Електронне фахове наукове видання «Кібербезпека: освіта, наука, техніка»*, 1(29). <https://doi.org/10.28925/2663-4023.2025.29.937>

**Serhii Nuzhnyi**

Candidate of Technical Sciences, Associate Professor,
Head of the Department of Computer Technologies and Information Security
Admiral Makarov National University of Shipbuilding, Mykolaiv, Ukraine
ORCID: 0000-0002-7706-0453
s.nuzhnyi@gmail.com

NEURAL NETWORK MODEL FOR ASSESSING THE SECURITY LEVEL OF HEAVILY NOISE-MASKED SPEECH INFORMATION BASED ON THE RII STRUCTURAL FRAMEWORK

Abstract. This study addresses the problem of determining the security level of speech information under complex acoustic and vibroacoustic interference, where traditional speech-quality metrics (SNR, STI, SII, PESQ, STOI) fail to reflect the actual capability of modern reconstruction algorithms (HMM, DNN, BayesianNN, GAN) to recover the semantic content of intercepted signals. Under low signal-to-noise ratios, and even after applying active vibroacoustic masking, a substantial portion of speech structures remains recoverable, creating a risk of information leakage. The absence of a quantitative criterion that unifies the physical, linguistic, and neural-network parameters of the signal and allows evaluation of the real possibility of semantic reconstruction defines the central scientific and technical challenge of this work. To address this, an integral model is proposed for estimating the residual speech information of heavily noise-masked speech, taking into account spectral structure, contextual linguistic dependencies, and the capabilities of modern reconstruction systems. The multi-level architecture of the model includes an analytical spectral description of the signal (A_i , μ_i , σ_i^* , $Z_o(f)$, $s(f)$), a Bayesian-Markov triphone structure, and a multilayer emission neural network built using CNN, MLP, and BayesianNN components. The spectral level provides a formal description of energy peaks and adaptive smoothing; the linguistic level reflects probabilistic regularities of transitions between triphones; the neural-network level integrates all feature types and models uncertainty in emission probabilities. Based on the synthesis of these levels, the residual intelligibility criterion (RII) is formulated, quantitatively characterizing the ability of a potential interceptor to recover message content after interference and filtering. A threshold value RII^* is defined, interpreted as a conditional boundary between information-dangerous and information-insufficient reconstruction regimes. The proposed model can be applied in technical information-protection systems, testing laboratories, and evaluation complexes for active vibroacoustic interference. The obtained results establish scientific and technical foundations for developing instrumental methods to determine the level of residual speech informativeness and enhance its protection.

Keywords: heavily noise-masked speech; Residual Intelligibility Index (RII); triphone model; emission neural network; BayesianNN.

REFERENCES (TRANSLATED AND TRANSLITERATED)

1. Deng, L., & Li, X. (2020). Deep learning in speech processing: A review. *IEEE Signal Processing Magazine*, 37(3), 107–139.
2. Kolossa, D., & Haeb-Umbach, R. (2018). *Robust Speech Recognition: A Probabilistic Approach*. Academic Press.
3. Pardo, J. M., et al. (2023). Bayesian deep learning for acoustic uncertainty modelling. *IEEE SPL*, 30, 642–646.
4. Zhao, Y., & Ma, J. (2020). Probabilistic reconstruction of masked speech using Bayesian neural networks. *Neural Networks*, 132, 229–241.
5. Arora, A., & Singh, R. (2021). Whisper-to-speech reconstruction using GANs. *Pattern Recognition Letters*, 152, 62–70. <https://doi.org/10.1016/j.patrec.2021.09.011>.
6. Koizumi, Y., et al. (2020). Speech enhancement using deep generative models. *IEEE/ACM TASLP*, 28, 1778–1788.



7. Pascual, S., Ravanelli, M., & Serra, J. (2020). SEGAN and its descendants: Generative enhancement of degraded speech. *IEEE SPS Magazine*, 37(6), 22–38.
8. Fogerty, D. (2017). Predicting speech intelligibility in noise. *Attention, Perception, & Psychophysics*, 79, 333–344.
9. Wang, D., & Chen, J. (2018). Supervised speech separation based on deep learning. *IEEE TASLP*, 26(10), 1872–1892.
10. Zolfaghari, M., & Sameti, H. (2023). End-to-end noise-robust ASR integrating spectral and linguistic features. *Computer Speech & Language*, 77, 101439.
11. ISO 22955:2021. Acoustics — Acoustic quality of open office spaces.
12. ISO 3382-3:2022. Acoustics — Measurement of room acoustic parameters — Part 3: Open public offices.
13. Reddy, C. K., et al. (2019). The Interspeech Deep Noise Suppression Challenge. *Interspeech*.
14. Li, A., et al. (2021). Two-stage deep enhancement for ultra-low SNR conditions. *Applied Acoustics*, 178, 108018.
15. Bhat, S., & Chatterjee, S. (2020). Complex-domain deep networks for speech enhancement at low SNR. *Speech Communication*, 122, 48–62.
16. Adiga, V. S., & Seltzer, M. L. (2020). End-to-end models for robust speech recognition in low-resource and noisy conditions. *ICASSP*.
17. Dong, L., & Xu, B. (2022). Noise-robust ASR using factorized time-domain convolutional networks. *Speech Communication*, 139, 15–27.
18. Sriram, J., & Chen, M. (2021). Transformer architectures for robust speech recognition. *ICASSP*.
19. Biswas, S., & Manocha, D. (2023). Transformer-based low-SNR speech enhancement using multi-scale spectral attention. *IEEE SPL*, 30, 124–128.
20. Hershey, J. R., et al. (2016). Deep clustering: Discriminative embeddings for segmentation and separation. *ICASSP*.
21. Qin, K., & Wang, D. (2020). Time-domain speech enhancement using deep learning: A survey. *Speech Communication*, 127, 1–16.
22. Andronic, E., & Dehak, N. (2019). End-to-end text-independent speaker verification for extremely noisy environments. *Interspeech 2019*, 4075–4079. <https://doi.org/10.21437/Interspeech.2019-2447>
23. Chen, J., Wang, X., & Xu, Y. (2021). Multi-task neural networks for joint speech enhancement and ASR. *IEEE TASLP*, 29, 2071–2084.
24. Kim, J., & Hahn, M. (2018). Voice conversion using GANs. *NeurIPS Workshop*.
25. Ma, J., & Zhao, Y. (2022). Audio-visual speech reconstruction in extremely noisy environments. *IEEE TPAMI*, 44(8), 4105–4117.
26. Alkin, B. (2021). *Applied acoustics research for speech privacy*. Springer. <https://doi.org/10.1007/978-3-030-67379-8>. ISBN 978-3-030-67378-1
27. Nuzhnyi, S. (2025). Udoskonalennia alhorytmu vidnovlennia linhvistychnoyi skladovoyi movnoyi informatsiyi fonemno-formantnym metodom dlia zadach otsinky rivnia yii zakhyshchenosti. Elektronne fakhove naukove vydannia *Kiberbezpeka: osvita, nauka, tekhnika*, 3(27). <https://doi.org/10.28925/2663-4023.2025.27>
28. Nuzhnyi, S. (2025). Adaptivnyi fonemno-spektralnyi metod vidnovlennia movlennia v umovakh aktyvnykh akustychnykh zavud. Elektronne fakhove naukove vydannia *Kiberbezpeka: osvita, nauka, tekhnika*, 4(28). <https://doi.org/10.28925/2663-4023.2025.28.794805>
29. Nuzhnyi, S. (2025). Otsiniuvannia rivnia zakhyshchenosti skladnozashumlendnoyi movnoyi informatsiyi za kryteriiem zalyshkovoyi rozbirlyvosti movy. Elektronne fakhove naukove vydannia *Kiberbezpeka: osvita, nauka, tekhnika*, 1(29). <https://doi.org/10.28925/2663-4023.2025.29.937>

