

Б. Ю. Варер**ТЕХНОЛОГІЯ ДЕКОМПОЗИЦІЇ ТА АНСАМБЛЮВАННЯ МОДЕЛЕЙ
ДЛЯ ЗМЕНШЕННЯ НЕВИЗНАЧЕНОСТІ ТА ПІДВИЩЕННЯ
УЗГОДЖЕНОСТІ ПРОГНОЗІВ НА ОСНОВІ LLM**

У статті запропоновано нову технологію прогнозування часових рядів із високим рівнем невизначеності, побудовану на основі ансамблю простих моделей та великої мовної моделі (LLM), яка визначає структурні параметри ансамблю, відповідно до теоретично обґрунтованого критерію оптимальності. Актуальність проблеми зумовлена складністю коротко- та середньострокового прогнозування даних екологічного моніторингу, епідемічних процесів, фінансових даних тощо після різких змін ситуації, коли треба робити прогноз за короткими рядами даних, за умов невизначеності, нерегулярності, неповноти та нестационарності вхідних даних.

Запропоновано теоретично обґрунтовану технологію зменшення прогнозової невизначеності та підвищення узгодженості прогнозів нестационарних часових рядів, яка базується на декомпозиції часових даних на локальні інтервали та оптимальному ансамблюванні моделей із використанням інтегрального критерію втрат. Зменшення прогнозової невизначеності трактується не як механічне звуження інтервалів прогнозування, а як зменшення некаліброваної та неінформативної невизначеності шляхом досягнення узгодженого балансу між точністю прогнозу, шириною та покриттям інтервалів прогнозування. Удосконалено підхід до ансамблевого прогнозування часових рядів шляхом формалізації задачі вибору структури ансамблю як задачі мінімізації інтегральних втрат, що одночасно враховують похибку прогнозу, ширину та каліброваність інтервалів прогнозування, а також структурну узгодженість розподілу ваг моделей, що дозволяє отримувати більш стабільні та узгоджені прогнозні рішення. Дістало подальшого розвитку системно-аналітичне обґрунтування використання простих локальних моделей на коротких часових інтервалах, за рахунок доведення умов, за яких ансамблювання таких моделей забезпечує зменшення прогнозової невизначеності та підвищення узгодженості результатів, що дозволяє застосовувати запропоновану технологію як базовий механізм координації рішень в мультиагентних системах та ансамблях LLM.

Проведено прикладне дослідження прогнозування COVID-19 в Україні у період першої великої хвилі в кінці 2020 року на основі історичних даних лише за півроку і за відсутності подібних рядів у минулому. На основі віконного аналізу (28 днів історії, 14 днів прогнозу) показано, що запропонована технологія забезпечує у 2,3 разів меншу метрику WAPE (аналог відносної похибки) порівняно з відомими моделями та забезпечує суттєво краще каліброване прогнозування, що проявляється у досягненні емпіричного покриття інтервалів на 0,85, за відсутності покриття в моделях-аналогах. Це свідчить не лише про зменшення похибки та некаліброваної прогнозової невизначеності, а й про підвищення узгодженості прогнозів, оскільки точкові та інтервальні оцінки формують єдиний сценарій динаміки процесу, який не містить протиріч.

Результати демонструють можливість використання гібридних LLM-керованих ансамблів для задач підтримки прийняття рішень у ситуаціях високої невизначеності та за відсутності структурно сталих моделей.

Ключові слова: часові ряди, ансамбль моделей, прогнозування, велика мовна модель, інтелектуальна модель, машинне навчання, невизначеність, ентропія, COVID-19.

Вступ

Сучасні підходи до прогнозування та підтримки прийняття рішень дедалі частіше ґрунтуються на використанні великих мовних моделей (LLM), які здатні працювати з неструктурованими даними, інтегрувати різномірну інформацію та формувати складні висновки. Водночас, численні дослідження, зокрема й авторські [1], показують, що результати роботи LLM є дуже чутливими до невизначеності вхідних даних, структури промптів та неповноти контексту, що призводить до зниження узгодженості відповідей і прогнозів. Особливо гостро ця проблема проявляється у задачах аналізу та прогнозування часових рядів реальних складних систем, де дані є нестационарними, фрагментованими,

містять пропуски та значну частку латентних факторів. За таких умов навіть використання кількох паралельних LLM або агентів на їх основі не гарантує узгодженості результатів без спеціальних механізмів зменшення невизначеності та координації рішень. Прикладами таких задач є прогнозування даних екологічного моніторингу, рядів курсу валют, даних про поширення захворювань, зокрема на коронавірус, тощо. Особливо складними є ситуації, коли режими суттєво та часто змінюються і прогноз слід робити на основі доволі коротких рядів спостережень. Наприклад, дуже складно було моделювати зростання кількості нових хворих на коронавірус під час першої «хвилі» у 2020 році в усіх країнах, зокрема й в Україні.

Подібними задачами займаються багато вчених у різних країнах. В останні роки все більша увага приділяється використанню LLM для підвищення ефективності певних операцій, за рахунок методів глибокого навчання, нових видів декомпозиції, моделей зі змішуванням з LLM у різних масштабах, використання семантичних просторів та ансамблювання різного роду моделей, але це спричиняє появу нових проблем, зокрема щодо узгодженості прогнозів, кількісного оцінювання ступеню невизначеності та обґрунтування оптимальності ансамблю за певним інтегральним критерієм [2 – 8]. Теоретично, моделі можуть надавати прогнози з певною доволі широкою зоною невизначеності, але це має мало практичного сенсу і потім, якщо їх подати на вхід LLM, тоді це погано впливає на узгодженість варіантів сценаріїв і відповідних рекомендацій на їх виході. Слід зазначити, що в межах цього дослідження узгодженість прогнозів розуміється як відсутність суперечностей між точковими та інтервальними оцінками прогнозів моделі, а у разі ансамблювання низки моделей, усі їх прогнози теж повинні узгоджуватись між собою. А зменшення прогнозової невизначеності розуміється не як мінімізація ширини інтервалів прогнозування, а як усунення некаліброваної невизначеності та підвищення узгодженості прогнозів, що є принципово важливим для задач підтримки прийняття рішень. Отже, важливо і підвищити точність прогнозування, і зменшити некалібровану прогнозову невизначеність, і збільшити узгодженість різних прогнозів заданих коротких рядів даних.

Метою дослідження є зменшення прогнозової некаліброваної невизначеності та підвищення узгодженості прогнозів нестационарних часових рядів за умов обмеженого обсягу даних.

Загальна постановка проблеми та ідея щодо побудови моделі

Прогнозування часових рядів реальних складних систем часто ускладнюється через малу кількість спостережень, нестационарність процесу з різними режимами динаміки, вплив багатьох зовнішніх факторів, які або не вимірюються взагалі, або вимірюються з великими похибками та затримками, нестабільність взаємозв'язків між змінними в часі тощо. За таких умов, побудова єдиної моделі для всього ряду або неможлива, або дає нестійкі та неточні прогнози. Тому доцільною альтернативою побудові єдиної моделі виступає декомпозиційний підхід, згідно якого замість однієї глобальної моделі будується набір локальних моделей для окремих інтервалів чи режимів з подальшим агрегуванням їх прогнозів.

Застосуємо функціонально-структурну декомпозицію моделей часових рядів, яка передбачає поєднання трьох підходів одночасно:

1. Часова декомпозиція: часовий ряд розбивається на вікна, що перекриваються, фіксованої чи адаптивної довжини, всередині яких припускається квазістационарність процесу.
2. Режимна декомпозиція: у кожному вікні допускається, що процес перебуває в одному з характерних режимів (зростання, спад, стабілізація тощо), для кожного з яких властиві певні закономірності.
3. Модельна декомпозиція: одного інтервалу будується набір локальних моделей, що враховують різні аспекти динаміки (такі як інерційність, чутливість, сезонність), замість однієї глобальної моделі.

Обмежимося тільки такими задачами, які допускають ефективне застосування такого виду

декомпозиції до даних і до моделі, яка будується для їх прогнозування. Крім того, сформулюємо ще ряд гіпотез і будемо вважати, що вони теж виконуються:

Гіпотеза Н1 (про локальну квазістаціонарність): існує така довжина вікна L , за якої часовий ряд всередині кожного вікна можна адекватно описати простими моделями.

Це дозволяє уникнути проблем глобальної нестационарності.

Гіпотеза Н2 (про скінченність режимів): динаміка процесу у кожному вікні належить до одного режиму зі скінченної множини типових режимів, які можуть повторюватись в часі.

Гіпотеза Н3 (про перевагу простих моделей): на коротких інтервалах прості моделі з невеликою кількістю параметрів дають стабільніші оцінки, ніж складні моделі з великою кількістю параметрів.

Гіпотеза Н4 (про відсутність універсальної моделі): жодна локальна модель не є оптимальною для всіх інтервалів і режимів динаміки.

Це обґрунтовує необхідність побудови оптимального рішення у вигляді ансамблю.

Лема 1.

За умови виконання гіпотез Н1-Н3, для достатньо коротких інтервалів L існує ансамбль простих локальних моделей, який забезпечує меншу сумарну похибку та меншу невизначеність прогнозу, ніж будь-яка окрема складна модель.

Доведення.

Цю лему легко довести, згадавши, що на коротких інтервалах дисперсія оцінок параметрів складних моделей зростає пропорційно їх параметричній розмірності. Відповідно до гіпотези Н3, прості моделі, хоч і менш універсальні, забезпечують стабільніші оцінки. Ансамблювання таких моделей дозволяє зменшити випадкову складову похибки та частково компенсувати їхні систематичні обмеження, що призводить до зменшення загальної похибки прогнозу.

Формалізація задачі

Припустимо, що задано цільовий часовий ряд y_t та, за наявності, пов'язані додаткові ряди $x_t \in \mathbb{R}^P$, $t = 1, \dots, T$.

Для прогнозування використовується ковзне вікно довжини L та горизонт прогнозу H . Для кожного моменту твизначимо:

Вектор спостережень (історія): $H_t = \{(y_t, x_t)_{t=\tau-L+1}^{\tau}\}$.

Вектор майбутніх значень (прогноз): $F_t = \{(y_t, x_t)_{t=\tau+1}^{\tau+H}\}$.

Задача полягає в тому, щоб для кожного τ побудувати прогноз траєкторії $\hat{y}_\tau = (\hat{y}_{\tau,1}, \dots, \hat{y}_{\tau,H})$ та інтервал невизначеності $[l_\tau, u_\tau]$.

Формалізуємо множину моделей-експертів та їх індикатори невизначеності.

Припустимо, що задано набір з N базових моделей:

$$M = \{m_1, \dots, m_N\}.$$

Кожна модель m_i за спостереженнями $y_{\tau-L+1:\tau}$ формує точковий прогноз $y_\tau^{(i)} \in \mathbb{R}^H$ та оцінку невизначеності у вигляді прогнозного інтервалу $[l_\tau^{(i)}, u_\tau^{(i)}]$. Для подальшого теоретичного аналізу зручно параметризувати невизначеність через дисперсію:

$$\sigma_\tau^{2(i)} \in \mathbb{R}_+^H.$$

Тоді для інтервалу $[l_\tau^{(i)}, u_\tau^{(i)}]$:

$$\begin{aligned} l_\tau^{(i)} &= \hat{y}_\tau^{(i)} - z\sigma_\tau^{(i)}, \\ u_\tau^{(i)} &= \hat{y}_\tau^{(i)} + z\sigma_\tau^{(i)}, \end{aligned}$$

де $z > 0$ – фіксований квантиль стандартного нормального розподілу (наприклад, 1,96).

Формалізуємо ансамбль моделей та його індикатори невизначеності.

Ансамбль задається вагами, що належать симплексу:

$$\Delta_N = \{w \in \mathbb{R}^N: w_i \geq 0, \sum w_i = 1\}, i = 1, \dots, N.$$

Для кожного τ прогноз ансамблю визначимо наступним чином:

– точковий прогноз:

$$\hat{y}_\tau(w) = \sum w_i \hat{y}_\tau^{(i)},$$

– агрегована дисперсія:

$$\sigma_\tau^2(w) = \sum w_i \left(\sigma_\tau^{2(i)} + \left(\hat{y}_\tau^{(i)} - \hat{y}_\tau(w) \right)^2 \right),$$

де $(\cdot)^2$ – покомпонентний квадрат.

Тоді прогнозний інтервал ансамблю матиме вигляд:

$$l_\tau(w) = \hat{y}_\tau(w) - z\sigma_\tau(w),$$

$$u_\tau(w) = \hat{y}_\tau(w) + z\sigma_\tau(w).$$

Формалізуємо інтегральний критерій, який буде одночасно забезпечувати:

- точність прогнозування шляхом мінімізації похибок MAE (з англ. "Mean Absolute Error" – середня абсолютна похибка), SMAPE (з англ. "Symmetric Mean Absolute Percentage Error" – симетрична середня абсолютна відсоткова похибка) тощо;
 - зменшення невизначеності шляхом зменшення ширини інтервалу невизначеності;
 - калібрування прогнозних інтервалів (штраф за недотримання цільового покриття);
 - узгодженість ваг (регуляризацію структури) шляхом штрафування за зростання ентропії.
- Для фіксованого моменту τ (одного вікна) розглянемо чотири складові, які разом визначають якість прогнозу ансамблю з вагами w :

1. Середня похибка точкового прогнозу, виміряна за допомогою робастної метрики такої, як, наприклад, MAE або SMAPE:

$$L_{err}(\tau, w) = \frac{1}{H} \sum |y_{\tau+h} - \hat{y}_{\tau,h}(w)|.$$

2. Міра невизначеності (ширина прогнозного інтервалу). Метою цієї складової є зменшення невизначеності шляхом мінімізації середньої ширини прогнозних інтервалів на горизонті прогнозу:

$$L_{wid}(\tau, w) = \frac{1}{H} \sum (u_{\tau,h}(w) - l_{\tau,h}(w)).$$

3. Калібрування (дотримання цільового покриття). Таким чином, інтервали оцінюються не лише за шириною, а й за покриттям відносно істинних значень. Для цього використовується емпіричне покриття – частка кроків $h = 1, \dots, H$ горизонту, для яких істинне значення $y_{\tau+h}$ належить прогнозному інтервалу:

$$L_{cov}(\tau, w) = \max\{0, c^* - \hat{c}(\tau, w)\},$$

$$\hat{c}(\tau, w) = \frac{1}{H} \sum 1 \{y_{\tau+h} \in [l_{\tau,h}(w), u_{\tau,h}(w)]\},$$

де $c^* \in (0,1)$ – цільове покриття, наприклад 0,9.

4. Узгодженість ваг ансамблю. Регуляризаційна складова контролює форму розподілу ваг w для забезпечення узгодженості ваг ансамблю:

$$R(w) = 1 - \frac{H(w)}{\log M},$$

де

$$H(w) = - \sum w_i \log(w_i + \varepsilon).$$

Оскільки інтегральний критерій мінімізується, то в нього включається:

$$L_{entH}(w) = 1 - R(w) = \frac{H(w)}{\log M}.$$

Припустимо, що T – набір вікон (тренувальних або верифікаційних). Визначимо інтегральну функцію втрат:

$$J(w) = \sum (\alpha L_{err}(\tau, w) + \beta L_{wid}(\tau, w) + \gamma L_{cov}(\tau, w) + \lambda L_{entH}(w)), \quad (1)$$

де $\alpha, \beta, \gamma, \lambda > 0$ – вагові коефіцієнти складових інтегрального критерію, можуть усі дорівнювати й 1.

Тоді задача полягає в оптимальному виборі структури:

$$w^* = \operatorname{argmin}_{w \in \Delta_M} J(w). \quad (2)$$

Сформулюємо теорему про існування оптимального ансамблю.

Теорема.

Припустимо, що виконуються гіпотези Н1–Н4, а також задано інтегральний критерій якості, що містить похибку прогнозу L_{err} , міру невизначеності L_{wid} та калібрування L_{cov} . Тоді існує оптимальний розподіл ваг w^* на симплексі Δ_N . Відповідний ансамбль є оптимальним серед усіх зважених комбінацій моделей.

Доведення. Оскільки симплекс Δ_N є замкнутою та обмеженою множиною в $\mathbb{R}^N$, він є компактим. Для фіксованих прогнозів $\{\hat{y}_\tau^{(i)}, \sigma_\tau^2(i)\}$ відображення $w \mapsto \hat{y}_\tau(w)$ є лінійним, а отже неперервним. Вираз для агрегованої дисперсії $\sigma_\tau^2(w)$ є сумою лінійних та квадратичних форм від w , тому також є неперервним, і, як наслідок, неперервними є $l_\tau(w)$ та $u_\tau(w)$.

Компоненти критерію якості L_{err} та L_{wid} є сумами композицій неперервних функцій (модуль, додавання, множення), тому – неперервні. Компонента L_{cov} містить індикаторну функцію, яка є розривною, водночас, у цьому випадку її внесок входить у вигляді штрафу $L_{cov}(\tau, w) = \max\{0, c^* - \hat{c}(\tau, w)\}$, тому L_{cov} є нижньо напівнеперервною на Δ_N . Таким чином, інтегральний критерій J є нижньо напівнеперервним на компактній множині Δ_N і, за теоремою Вейерштраса для нижньо напівнеперервних функцій на компактi існує точка мінімуму $w^* \in \Delta_N$.

Алгоритм застосування технології

Вхідні дані: ряди y_t, x_t ; довжини L, H ; набір моделей M ; коефіцієнти інтегрального критерію $\alpha, \beta, \gamma, \lambda > 0$; квантиль z ; цільове покриття c^* .

Вихідні результати: для кожного τ : $\hat{y}_\tau, [l_\tau, u_\tau]$, оцінка режиму/стану (за потреби).

Пропонується такий алгоритм застосування технології:

1. Провести декомпозицію і сформувані вікна: побудувати множину T ковзних вікон H_τ .
2. Побудувати ознаки: для кожного τ обчислити вектор ознак $\Phi(H_\tau)$ (тренд, волатильність, рівень, недавні помилки моделей тощо).
3. Прогнозування за моделями та визначення їх прогнозного інтервалу: для кожного $m_i \in M$ і $\tau \in T$ отримати $\hat{y}_\tau^{(i)}$ та $\sigma_\tau^{2(i)}$ (або $[l_\tau^{(i)}, u_\tau^{(i)}]$).
4. Сформувані структуру ансамблю (за допомогою LLM або, у першому наближенні – з використанням іншого оптимізатору): використовуючи ознаки $\Phi(H_\tau)$, сформувати кандидати ваг $w_\tau \in \Delta_N$.
5. Агрегувати прогноз та інтервал: обчислити $\hat{y}_\tau(w_\tau), l_\tau(w_\tau), u_\tau(w_\tau)$.
6. Визначення запобіжників (допустимих перетворень ваг): застосувати оператор Π (проекція на допустиму множину/правила), що гарантує: невід'ємність і нормування ваг, обмеження дрейфу від «базової» моделі, обмеження на сумарну масу прогнозу, формально: $w_\tau = \Pi(w_\tau) \in \Delta_N$.
7. Оцінювати якість на T та мінімізувати J : обчислити J на train-вікнах і уточнити процедуру побудови w_τ (оновлення правил/промпта), щоб зменшити J .

Приклад застосування технології для задачі прогнозування щоденного приросту кількості хворих на коронавірус в Україні у 2020 році

Перевіримо ефективність розробленої технології побудови ансамблю простих моделей для прогнозування реальних даних з тренувальним датасетом малої тривалості (менше року) на прикладі ряду даних про щоденні прирости кількості нових підтверджених хворих на коронавірус в Україні під час першої великої «хвилі» в кінці 2020 року.

При Президії НАН України була створена міжвідомча Робоча група з математичного моделювання проблем, пов'язаних з епідемією коронавірусу SARS-CoV-2 в Україні (базова установа – Інститут проблем математичних машин і систем (ІПММС) НАН України), яка протягом 2020–2022 років кожні 1–2 тижні робила прогнози, готувала аналітику, оформляла у вигляді звітів «Прогноз РГ-XX» (XX – порядковий номер, наприклад «Прогноз РГ-31» [9])

та пересилала в РНБО України та в основні урядові установи, які координували формування та впровадження управлінських рішень щодо запобігання процесу поширення, мінімізації шкоди та розроблення методик лікування від коронавірусу в Україні. Наразі, усі звіти доступні на сайті Президії НАН України [10]. У 25 з 62 звітів, які відповідали періодам значних «хвиль» щодо приростів щоденної кількості нових хворих, присутній прогноз двома моделями – за допомогою моделі SEIR-U команди з НАН України під керівництвом заступника директора ІПММОН НАН України В. Бровченка та команди з ВНТУ під керівництвом завідувача кафедри САІТ ВНТУ В. Мокіна [9]. Основна інформація про обидві моделі наведена у статті [11]. А у статті [12] А. Лосенко – один із членів команди В. Мокіна – навів аналіз точності прогнозування обома моделями в усіх 25 звітах. Причому, наведена як точність на валідаційних даних (2 останні тижні до прогнозу), так і тестові дані (2 наступні тижні після прогнозу). В якості метрики була використана WAPE (англ. «Weighted Absolute Percentage Error»), помножена на 100 % [12, 13]:

$$WAPE(\%) = 100 \cdot \frac{\sum_{i=1}^n |y_i - \hat{y}_i|}{\sum_{i=1}^n |y_i|}.$$

Іноді її помилково називають усередненою за інтервал відносною похибкою, що, строго математично, не дуже коректно, але, за певних умов ці значення є подібними.

Щодо звіту «Прогноз РГ-30» у таблиці на сторінці 57 статті [12] видно, що станом на 30.11.2020 р. за прогнозами зробленими на 1.12.2020 – 13.12.2020 рр. похибка за моделлю на основі Prophet склала 32,47 % (похибка на валідаційному датасеті була 3,2 %), а за моделлю SEIR-U – 30,81 %. Причини цього легко зрозуміти, якщо проаналізувати графіки у наступному звіті «Прогноз РГ-31» (рис. 1).

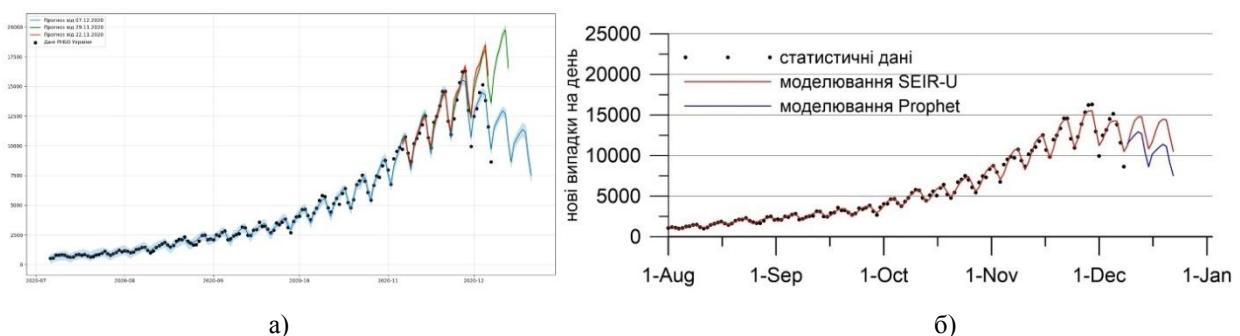


Рис. 1. Графіки первинних даних (чорні крапки на обох графіках) про щоденні прирости кількості нових підтверджених хворих на коронавірус в Україні під час першої великої «хвилі» у другій половині 2020 року та прогнозів за моделлю на основі Prophet (а) та за нею ж та моделлю SEIR-U (б) [9]

Як видно з рис. 1а (рис. 18 звіту [9]), прогноз зеленою лінією, зроблений в момент, що передував піку хвилі, був напрочуд неточним, хоча волатильність була передбачена доволі точно, але – не середнє значення. А як видно з рис. 1б (рис. 24 того ж звіту [9]), прогноз червоною лінією моделлю SEIR-U, який, на відміну моделі Prophet, враховував багато різних ознак, зміг краще спрогнозувати вихід на плато, але не зміг спрогнозувати різкий спад і не дуже добре спрогнозував волатильність. Тому обидві моделі показали таку низьку точність. Одним з важливих факторів було те, що така динаміка в Україні мала місце вперше і не була достатньо досліджена. Однак, той факт, що модель SEIR-U спромоглась спрогнозувати відсутність зростання, означає, що дані містили необхідну для цього інформацію, яку можна спробувати обробити більш ефективно. Саме для цього варіанту даних і була застосована розроблена у цій статті технологія.

Для кожного прогнозного вікна було побудовано кілька простих моделей, таких як: швидка модель тренду (оцінка тенденції за останні 7 днів), базова модель тренду (за останні 14 днів), стабільна модель тренду (за останні 21 день), а також тижнева найважливіша модель, що відтворює тижневу сезонність (повторення шаблону попереднього тижня). Для

кожної моделі було отримано точковий прогноз і прогнозний інтервал. Ваги ансамблю визначалися великою мовною моделлю (LLM) для кожного прогнозного вікна. На етапі тренування застосовувалася парадигма «вчитель-учень» на навчальній частині часових вікон (історичні дані до тестового інтервалу): додаткова LLM (GPT-4.1) генерувала правила-корекції (lessons), які потім використовувалися основною LLM (Llama-3.1). На основі цих ваг було обчислено прогноз ансамблю як зважену суму середніх прогнозів моделей.

На рис. 2 наведено значення складових інтегрального критерію (J): L_{err} – нормалізована похибка прогнозу, L_{wid} – нормалізована середня ширина довірчого інтервалу, cov – емпіричне покриття інтервалом, L_{cov} – штраф за недопокриття, L_{ent} – ентропійна регуляризація ваг (лише для цього методу), а $J_{partial} = L_{err} + L_{wid} + L_{cov}$ та J – відповідно частковий і повний критерій (коефіцієнти $\alpha, \beta, \gamma, \lambda = 1$).

[1]: Складові інтегрального критерію (2020-12-01 ... 2020-12-13) для запропонованої технології і Prophet

	Метод	L_{err}	L_{wid}	cov	L_{cov}	L_{ent}	$J_{partial}$	J_{full}
0	Запроп. техн.	0.1348	0.5695	0.8462	0.0000	0.0739	0.7043	0.7782
1	Prophet	0.3159	0.0486	0.0000	0.8000	nan	1.1645	1.1645

Рис. 2. Значення складових інтегрального критерію J для запропонованого методу та моделі Prophet на інтервалі 2020-12-01–2020-12-13 pp.

Як видно на рис. 2, для розглянутого інтервалу запропонований метод забезпечує достатнє покриття ($cov = 0,85$) при меншій похибці, тоді як для Prophet-інтервал не накриває жодне істинне значення ($cov = 0$), що призводить до суттєвого зростання $J_{partial}$. Отже, інтегральний критерій узгоджено відображає перевагу запропонованого підходу за балансом «точність-невизначеність». На рис. 3 наведено графіки результатів прогнозування, разом зі значеннями прогнозів моделей-аналогів Робочої групи – моделі на основі Prophetта моделі SEIR-U зі звіту «Прогноз РГ-30» [14].

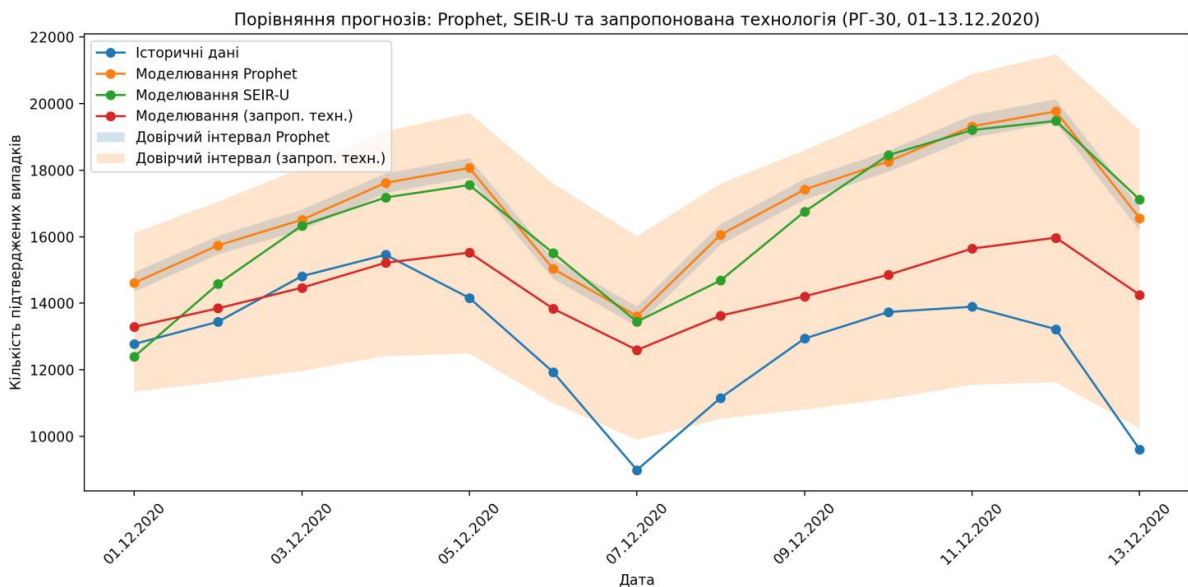


Рис. 3. Порівняння прогнозів кількості нових підтверджених випадків COVID-19 в Україні (РГ-30, 01–13.12.2020): історичні дані, прогнози моделей Prophet і SEIR-U та прогноз запропонованої технології; для Prophet і запропонованої технології наведено довірчі (прогнозні) інтервали

Як видно з рис. 3, запропонована технологія в середньому краще узгоджується з історичними даними, тоді як модель на основі Prophet і модель SEIR-U переоцінюють значення показника на більшості днів цього інтервалу. Отриманий за запропонованою технологією прогноз також характеризується підвищеною узгодженістю, яка проявляється у відсутності суперечностей між точковим прогнозом та інтервалом прогнозування, а також у формуванні єдиного узгодженого сценарію розвитку процесу без взаємовиключних тенденцій.

Отримане значення метрики WAPE 13,48 % є більш, ніж у 2,3 разів меншим за це ж значення щодо моделей на основі Prophet (32,47 %) та за моделлю SEIR-U (30,81 %), що продемонструвало ефективність запропонованої технології. Таким чином, на розглянутому інтервалі запропонована технологія одночасно забезпечує і зменшення відносної похибки прогнозу, і усунення некаліброваної прогнозової невизначеності, оскільки прогнозні інтервали охоплюють більшість істинних значень, на відміну від моделей-аналогів.

Висновки

У статті розглянуто задачу прогнозування нестационарних часових рядів за умов обмеженого обсягу даних і підвищеної невизначеності, що є характерним для широкого класу задач системного аналізу та підтримки прийняття рішень. Основну увагу зосереджено на проблемі підвищення узгодженості прогнозів, отриманих від множини паралельно функціонуючих моделей або агентів.

Запропоновано здійснювати декомпозицію часових даних на локальні інтервали та побудові для кожного з них множини простих моделей. Показано, що такий підхід дозволяє врахувати зміну режимів динаміки та частково компенсувати нестационарність процесу, не вимагаючи значних обсягів навчальних даних. Введено інтегральний критерій якості прогнозування, який одночасно враховує точність прогнозу, міру прогнозової невизначеності, каліброваність прогнозних інтервалів та структурну узгодженість ансамблю моделей. Теоретично доведено існування оптимальної структури ансамблю локальних моделей за заданим інтегральним критерієм. Сформульовано та доведено лему, яка пояснює перевагу ансамблю простих локальних моделей на коротких часових інтервалах і показує, що за обмеженого обсягу даних оптимальне агрегування дозволяє зменшити як похибку прогнозу, так і прогнозну невизначеність.

В цілому, вперше запропоновано теоретично обґрунтовану технологію зменшення прогнозової некаліброваної невизначеності та підвищення узгодженості прогнозів нестационарних часових рядів, що базується на декомпозиції часових даних на локальні інтервали та оптимальному ансамблюванні моделей із використанням інтегрального критерію втрат. Отримані результати підтверджують, що підвищення якості прогнозування досягається не за рахунок штучного зменшення ширини інтервалів прогнозування, а шляхом підвищення їх каліброваності та узгодженості з реальною динамікою процесу, що забезпечує зменшення некорисної прогнозової невизначеності.

Наведений приклад застосування підтвердив практичну реалізованість теоретичних положень і продемонстрував ефективність підходу в умовах коротких нестационарних часових рядів. Зокрема, приклад прогнозування на інтервалі 01.12.2020 – 13.12.2020 показав, що застосування запропонованої технології дозволило зменшити значення метрики WAPE до 13,48 %, тоді як для моделей Prophet та SEIR-U відповідні значення становили 32,47 % та 30,81 %, одночасно було досягнуто збільшення емпіричного покриття прогнозних інтервалів на 0,85, оскільки у моделей-аналогів покриття було нульовим. Це свідчить не лише про зменшення похибки та некаліброваної прогнозової невизначеності, а й про підвищення узгодженості прогнозів, оскільки точкові та інтервальні оцінки формують єдиний сценарій динаміки процесу, який не містить протиріч.

СПИСОК ЛІТЕРАТУРИ

1. Варер Б. Ю., Мокін В. Б. Аналіз впливу інформаційної невизначеності на узгодженість рекомендацій великих мовних моделей для кризових сценаріїв, С. 53–55. Колективна монографія за матеріалами XXIV Міжнародної науково-практичної конференції «Інформаційно-комунікаційні технології для стійкості та відновлення» (Київ, 11-12 листопада 2025 р.) / За заг. ред. С. О. Довгого. К.: ТОВ «Видавництво «Юстон», 2025, 221 с. URL: https://itgip.org/wp-content/uploads/2025/11/1_zbirka_2025_dlja-sajtu.pdf.
2. Wang L., Dong K., Zhao X. Anovel LLM time series forecasting method based on integer-decimal decomposition. *Sci Rep.* 2025. Vol. 15, Art. 23004. DOI: 10.1038/s41598-025-06581-x.
3. Deep learning for time series forecasting: a survey / X. Kong et al. *Int. J. Mach. Learn. & Cyber.* 2025. Vol. 16. P. 5079–5112. DOI: 10.1007/s13042-025-02560-w.
4. Are Language Models Actually Useful for Time Series Forecasting? / M. Tan et al. *Advances in Neural Information Processing Systems.* 2024. Vol. 37. P. 60162–60191. DOI: 10.52202/079017-1922.
5. Time CMA: Towards LLM-Empowered Multivariate Time Series Forecasting via Cross-Modality Alignment / C. Liu et al. *Proceedings of the AAAI Conference on Artificial Intelligence.* 2025. Vol. 39, №18. P. 18780–18788. DOI: 10.1609/aaai.v39i18.34067.
6. LLM-Mixer: Multiscale Mixing in Large Language Models for Time Series Forecasting / M. Kowsher et al. *Proceedings of the 4th Table Representation Learning Workshop.* 2025. P. 156–165. DOI: 10.18653/v1/2025.trl-1.12.
7. S²IP-LLM: Semantic Space Informed Prompt Learning with Large Language Models for Time Series Forecasting / Z. Pan et al. *Proceedings of the 41st International Conference on Machine Learning.* 2024. Vol. 235. P. 39135–39153.
8. Deep learning-based time series forecasting / X. Song et al. *Artif. Intell. Rev.* 2025. Vol. 58. Article number 23. DOI: 10.1007/s10462-024-10989-8.
9. Прогноз розвитку епідемії COVID-19 в Україні на 7–21 грудня 2020 року («Прогноз РГ-31»). URL: <https://old.nas.gov.ua/UA/Messages/Pages/View.aspx?MessageID=7238>, дата звернення: лист. 2025.
10. Сайт Президії Національної академії наук України (старий сайт). Розділ "Протидія COVID-19". Прогноз розвитку епідемії COVID-19 в Україні. URL: <https://old.nas.gov.ua/UA/Activity/covid/Pages/wg.aspx>, дата звернення: лист. 2025.
11. Мокін В., Лосенко А., Яцолт А. Інформаційна технологія аналізу та прогнозування кількості нових випадків на коронавірус SARS-CoV-2 в Україні на основі моделі Prophet. *Вісник Вінницького політехнічного інституту.* 2020. №5. С. 71–83. <https://doi.org/10.31649/1997-9266-2020-152-5-71-83>.
12. Лосенко А. Інформаційна технологія прогнозування часового ряду кількості хворих на коронавірус на основі моделі Facebook Prophet. *Вісник Вінницького політехнічного інституту.* 2023. №5, С. 50–59. URL: <https://doi.org/10.31649/1997-9266-2023-170-5-50-59>.
13. Mokin V., Losenko A. COVID in UA: Prophetwith 4, Nd seasonality, Kaggle Notebook. URL: <https://www.kaggle.com/code/vbmokin/covid-in-ua-prophet-with-4-nd-seasonality>. Accessed: Nov. 2025.
14. Прогноз розвитку епідемії COVID-19 в Україні на 1–14 грудня 2020 року («Прогноз РГ-30»). URL: <https://old.nas.gov.ua/UA/Messages/Pages/View.aspx?MessageID=7215>, дата звернення: лист. 2025.

Стаття надійшла до редакції 15.12.2025.

Стаття пройшла рецензування 20.12.2025.

Варер Борис Юхимович – аспірант кафедри системного аналізу та інформаційних технологій, e-mail: androbor17@gmail.com.

Вінницький національний технічний університет.