

література



Навчально-методична

Міністерство освіти і науки України
Тернопільський національний технічний університет
ім. Івана Пулюя
Кафедра комп'ютерно-інтегрованих технологій

Методичні вказівки
для виконання лабораторних робіт
з дисципліни
«КОМП'ЮТЕРНІ СИСТЕМИ ШТУЧНОГО
ІНТЕЛЕКТУ»
для студентів спеціальності G7 «Автоматизація,
комп'ютерно-інтегровані технології та робототехніка»
денної та заочної форми здобуття освіти

Тернопіль – 2025

Методичні вказівки для виконання лабораторних робіт з дисципліни «Комп'ютерні системи штучного інтелекту» для студентів спеціальності G7 «Автоматизація, комп'ютерно-інтегровані технології та робототехніка» денної та заочної форми здобуття освіти / уклад. І. С. Дідич. // ТНТУ. – 2025. – С. 36.

Укладач: доктор філософії, доц., Ірина ДІДИЧ
Рецензент: д.т.н., проф., Олег ЯСНІЙ

Відповідальний за випуск: доктор філософії, доц., Ірина ДІДИЧ

Методичні вказівки для виконання лабораторних робіт з дисципліни «Комп'ютерні системи штучного інтелекту» для студентів спеціальності G7 «Автоматизація, комп'ютерно-інтегровані технології та робототехніка» денної та заочної форми здобуття освіти розглянуто і схвалено на засіданні кафедри комп'ютерно-інтегрованих технологій
Протокол № 1 від “ 28 ” серпня 2025 року.

Методичні вказівки для виконання лабораторних робіт з дисципліни «Комп'ютерні системи штучного інтелекту» для студентів спеціальності G7 «Автоматизація, комп'ютерно-інтегровані технології та робототехніка» денної та заочної форми здобуття освіти схвалено та рекомендовано до друку науково-методичною комісією факультету прикладних інформаційних технологій та електроінженерії
Протокол № 1 від « 29 » серпня 2025 року

ЗМІСТ

Лабораторна робота №1. Застосування методів підсилених дерев та випадкових лісів.....	5
Лабораторна робота №2. Застосування методів k-найближчих сусідів та опорно-векторних машин.....	10
Лабораторна робота №3. Застосування методу нейронних мереж до вирішення задач регресії з використанням навчального алгоритму зворотного поширення помилок і інструментальних засобів пакету STATISTICA.....	14
Лабораторна робота №4. Застосування методу нейронних мереж до вирішення задач класифікації з використанням навчального алгоритму зворотного поширення помилок і інструментальних засобів пакету STATISTICA.....	20
Лабораторна робота №5. Генерація зображень на основі дифузійних мереж....	24
Лабораторна робота №6. Використання мереж DALL-E 3, Midjourney, ImageFX, Adobe Firefly, Leonardo для генерації зображень.....	27
РЕКОМЕНДОВАНА ЛІТЕРАТУРА.....	34

ЛАБОРАТОРНА РОБОТА №1

Тема: Застосування методів підсилених дерев та випадкових лісів у пакеті STATISTICA.

Мета роботи:

Отримати практичні навички використання модулів Boosted Trees та Random Forest у середовищі STATISTICA, навчитися будувати, налаштовувати та оцінювати ансамблеві моделі машинного навчання.

Короткі теоретичні відомості:

Алгоритм підсилених дерев ґрунтується на рекурсивному розбитті вхідної множини на підмножини, асоційовані із класами. Алгоритм конструювання під-силених дерев складається з етапів побудови і скорочення дерев. Створюючи дерева, вибирають критерій розщеплення та зупинки навчання. Тоді як під час їх скорочення відсікають деякі гілки. Ідея методу випадкових лісів полягає у побудові ансамбля з певної кількості дерев прийняття рішень, які навчаються неза-лежно одне від одного. Підсумковий результат приймають після голосування всіх дерев ансамбля. Перевагами цього методу є висока точність прогнозування та здатність ефективно обробляти дані

Крок 1. Запуск системи STATISTICA

1. Відкрийте STATISTICA 12 через меню Пуск або піктограму на робочому столі.
2. Після першого запуску оберіть зручний інтерфейс (рис. 1.1): Ribbon Bar (Стрічка) або Classic Menu (Класичне).
3. Для зміни інтерфейсу використовуйте Options → Ribbon Bar або View → Ribbon Bar.

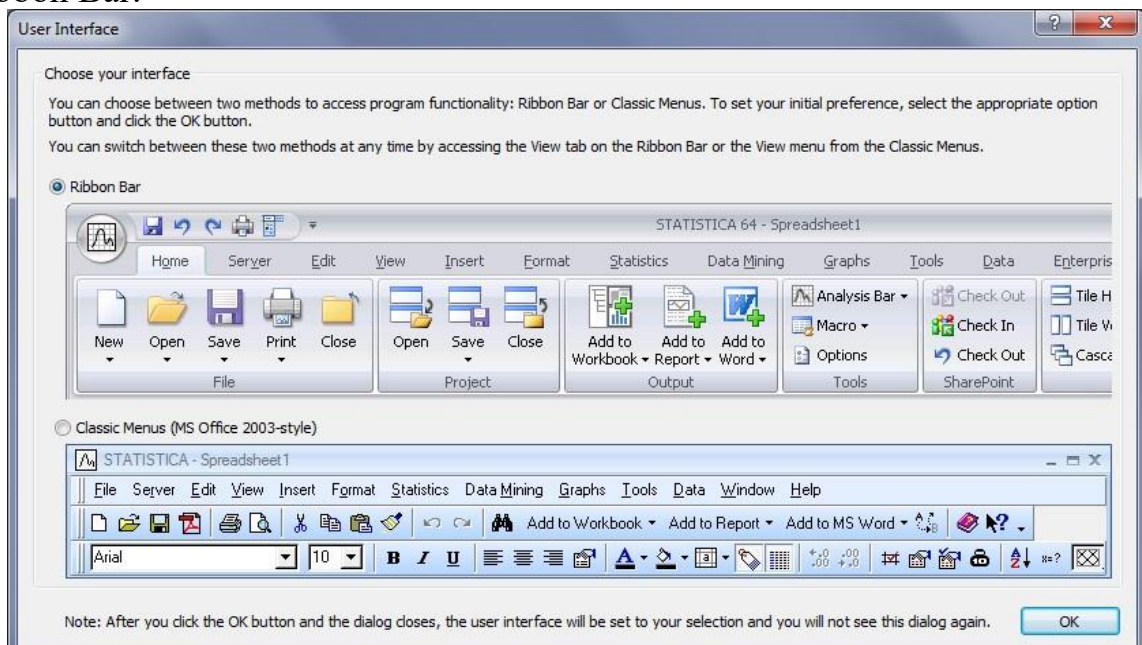


Рис. 1.1 – Вибір інтерфейсу при першому запуску STATISTICA.

Крок 2. Створення нового файлу даних

1. Виберіть File → New → Spreadsheet (Таблиця).
2. У діалоговому вікні задайте Number of variables, Number of cases, Display format.

3. Натисніть ОК, щоб створити порожню таблицю даних.

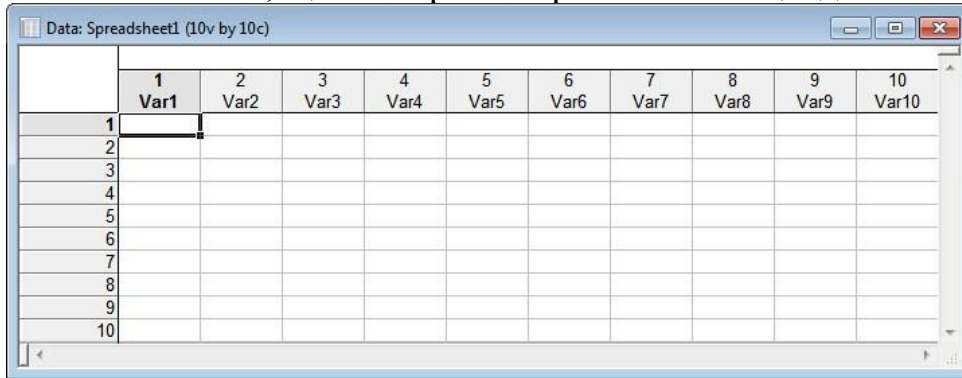


Рис. 1.2 – Створення нової таблиці даних у STATISTICA.

Крок 3. Введення та редагування змінних

1. Двічі клацніть у заголовку стовпця, щоб задати ім'я змінної.
2. Вкажіть Name, Type, Measurement Type, Label.
3. Для детального налаштування скористайтесь вкладкою Data → Variables → Specs.

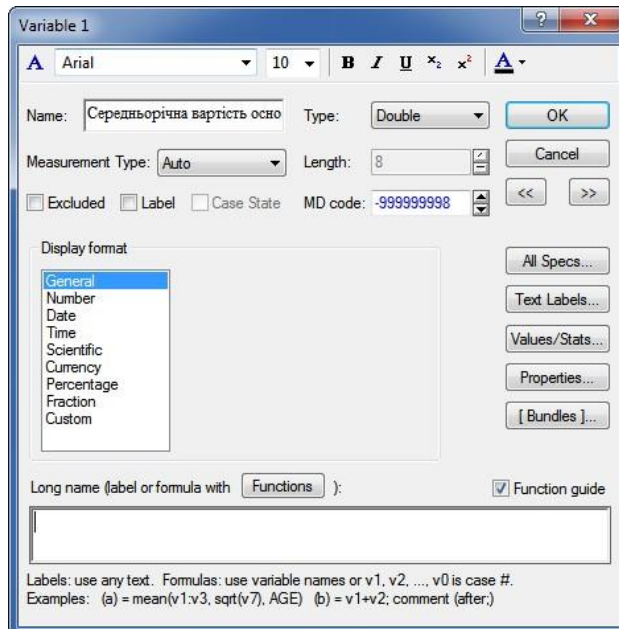


Рис. 1.3 – Діалогове вікно налаштування параметрів змінних.

Крок 4. Імпорт експериментальних даних

1. Виберіть File → Open → Data File.
2. Оберіть файл формату .STA, .XLS(X) або .CSV.
3. Перевірте коректність назв змінних і одиниць вимірювання.
4. Збережіть файл у форматі STATISTICA Spreadsheet (.sta).

	1	2	3	4	5	6	7	8	9	10	
	Log10(K)	Log10(V)	R	Var4	Var5	Var6	Var7	Var8	Var9	Var10	
1	0,3579	-9,27	0,2								
2	0,361	-10,451	0,2								
3	0,3673	-9,7496	0,2								
4	0,3754	-9,3354	0,2								
5	0,3945	-9,0057	0,2								
6	0,3972	-9,0209	0,2								
7	0,4242	-8,9355	0,2								
8	0,4346	-8,983	0,2								
9	0,4396	-8,6635	0,2								
10	0,493	-8,4962	0,2								
11	0,4953	-8,6556	0,2								
12	0,5201	-8,6216	0,2								
13	0,5222	-8,4789	0,2								
14	0,5534	-8,4828	0,2								
15	0,5933	-8,3799	0,2								
16	0,5942	-8,251	0,2								
17	0,5959	-8,4123	0,2								
18	0,5977	-8,5171	0,2								
19	0,5986	-8,5229	0,2								
20	0,6012	-8,301	0,2								

Рис. 1.4 – Вигляд електронної таблиці з імпортованими даними.

Крок 5. Побудова моделі Boosted Trees

1. Виберіть Data Mining → Boosted Trees → Regression Analysis.
2. Вкажіть Dependent variable (наприклад, da_dN) і Covariates (наприклад, DeltaK, R).
3. У вкладці Advanced задайте Number of Trees, Learning Rate, Tree Depth.
4. Натисніть Train для запуску навчання моделі.

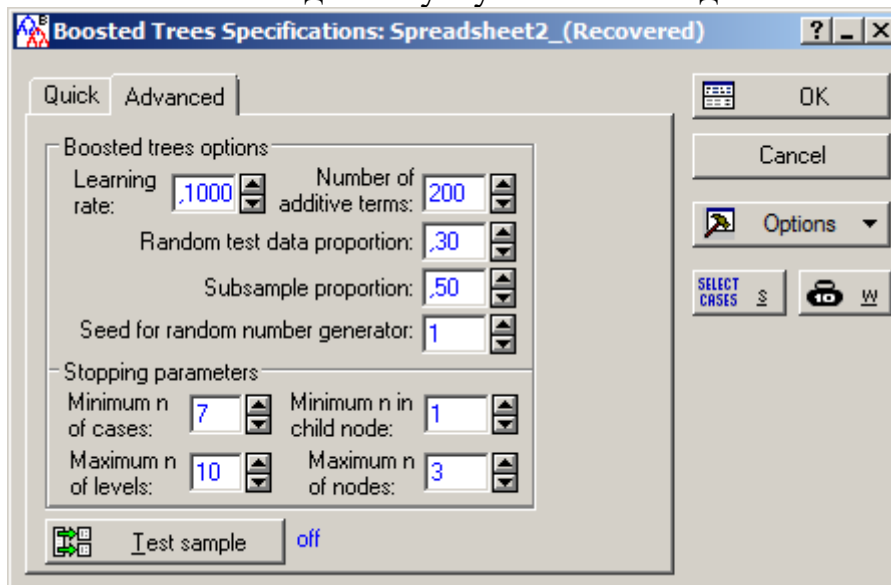


Рис. 1.5 – Вікно налаштування параметрів Boosted Trees.

5. У Results перегляньте похибку, графіки та важливість змінних.

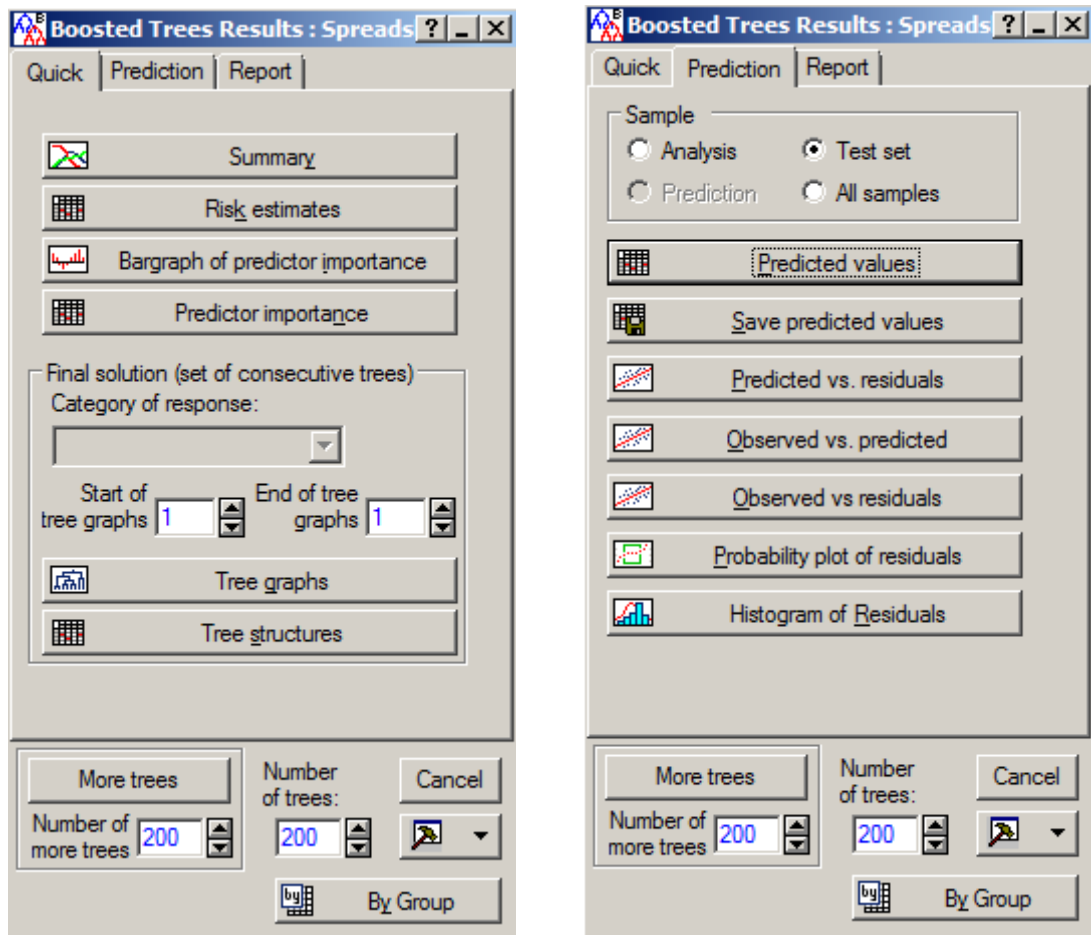


Рис. 1.6 – Приклад результатів моделі Boosted Trees.

Крок 6. Побудова моделі Random Forest

1. Виберіть Data Mining → Random Forests → Regression Analysis.
2. Вкажіть Dependent variable і Covariates.
3. У вкладці Advanced задайте Number of Trees, Variables per Split, Maximum Depth.
4. Натисніть Train, після завершення перегляньте вкладку Results.
5. Проаналізуйте Variable Importance – вплив кожного фактора.

Крок 7. Порівняння моделей

1. Перегляньте MSE (Mean Squared Error) та R^2 (Coefficient of Determination).
2. Порівняйте точність моделей Boosted Trees та Random Forest.
3. Зробіть висновок, яка модель має кращу узагальнюючу здатність.

Зміст звіту:

1. Тема, мета і завдання лабораторної роботи.
2. Короткі теоретичні відомості.
3. Опис використаних даних.
4. Етапи побудови моделей у STATISTICA.
5. Отримані результати (таблиці, графіки, похибки).
6. Висновки та рекомендації.

Контрольні запитання:

1. Що таке ансамблеві методи машинного навчання?
2. У чому полягає різниця між Boosting і Bagging?
3. Як обирається кількість дерев у Random Forest?
4. Які параметри впливають на швидкість навчання Boosted Trees?
5. Як визначити важливість змінних у STATISTICA?
6. Які показники використовуються для оцінки точності моделі?
7. Як уникнути перенавчання при побудові ансамблевих моделей?

ЛАБОРАТОРНА РОБОТА №2

Тема: Застосування методів k -найближчих сусідів та опорно-векторних машин у пакеті STATISTICA.

Мета роботи:

Закріпити теоретичні знання щодо принципів функціонування методів k -найближчих сусідів і опорно-векторних машин; засвоїти методику побудови, навчання, оцінювання і застосування цих методів для задач регресії та класифікації у STATISTICA 12.

Короткі теоретичні відомості:

Метод k -найближчих сусідів (k -NN) визначає належність нового об'єкта до класу на основі відстані до найближчих сусідів. Основними метриками є Евклідова, Манхеттенська, Мінковського.

Після обчислення відстаней вибираються k найближчих сусідів до точки, яку потрібно класифікувати чи для якої необхідно передбачити числове значення. Для класифікації об'єкт відноситься до класу, який найчастіше зустрічається серед k сусідів (метод більшості голосів). Для регресії прогнозоване значення визначається як середнє (або зважене середнє) значень сусідів.

Перевагами методу є простота реалізації та відсутність припущень про форму залежності між змінними.

Метод опорно-векторних машин (SVM) належить до алгоритмів, призначених для пошуку оптимальної межі між класами або поверхні, що найкраще апроксимує дані у випадку регресії. У випадку нелінійного поділу використовується ядрове перетворення (Kernel): лінійне, поліноміальне, RBF або сигмоїдне ядро. Після навчання якість моделі оцінюється за середньоквадратичною похибкою (MSE) або коефіцієнтом детермінації (R^2).

Крок 1. Підготовка даних

1. Запустіть STATISTICA 12.
2. Імпортуйте таблицю даних через File → Open → Data File.
3. Перевірте назви змінних: ΔK , R, V (da/dN).
4. Збережіть таблицю у форматі .sta для подальшого аналізу.

Крок 2. Побудова моделі методом SVM

1. Виберіть Data Mining → Machine Learning → Support Vector Machine.
2. У діалоговому вікні натисніть Variables.
3. Встановіть: Dependent variable – V (da/dN); Covariates – ΔK , R.
4. У вкладках Sampling, Kernels, Training задайте параметри: Kernel – RBF або Linear; C (Penalty) – 1.0; γ (Gamma) – 0.1–0.5 (рис.2.1).
5. Натисніть Train, після навчання відкрийте вкладку Predictions.
6. Порівняйте експериментальні та прогнозовані значення.

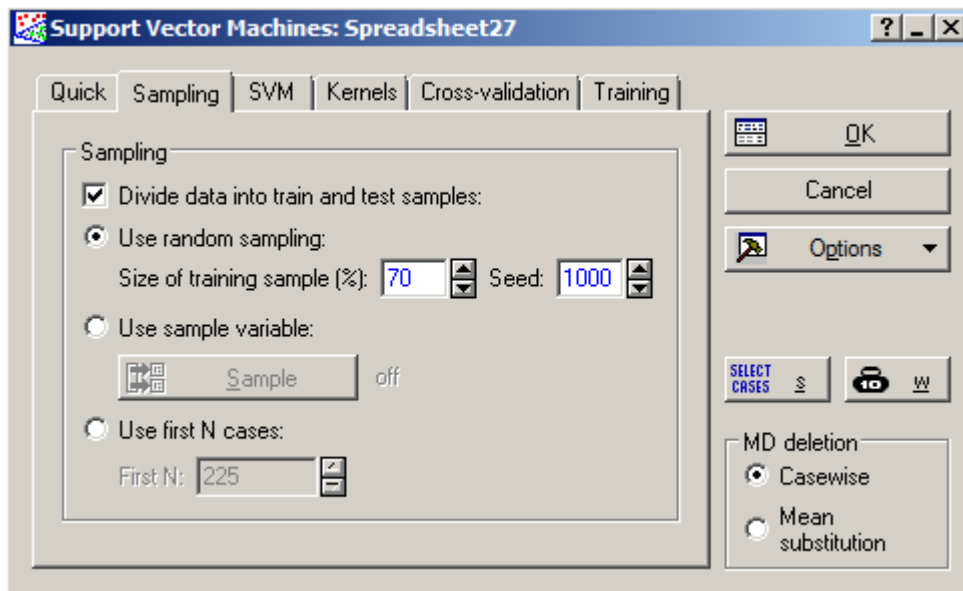


Рис. 2.1 – Налаштування вікна методу опорно-векторних машин у STATISTICA.

Крок 3. Аналіз результатів SVM

1. Перейдіть у вкладку Predictions (рис.2.2).
2. Перевірте стовпці: Experimental, Predicted, Residual.
3. Створіть стовпець Squared Residual для обчислення MSE.
4. Побудуйте графік Predicted vs Experimental для оцінки якості моделі.

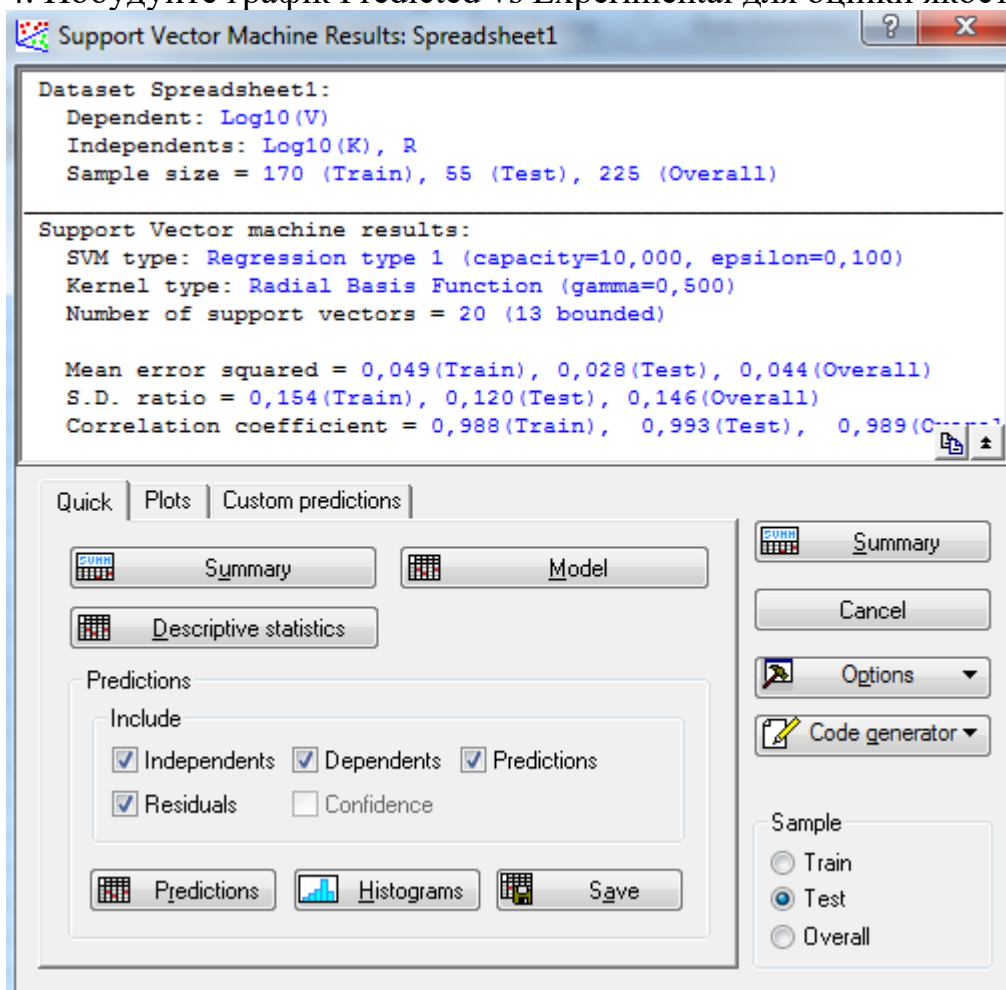


Рис. 2.2 – Результати прогнозування методом SVM.

Крок 4. Побудова моделі методом k -NN

1. Виберіть Data Mining → Machine Learning → K-Nearest Neighbors.
2. У меню Variables задайте: Dependent variable – V (da/dN); Covariates – ΔK , R .
3. У вкладці Training вкажіть Number of Neighbors (k) = 3–7.
4. Натисніть Train, потім перейдіть до Predictions (рис.2.3).
5. Оцініть точність прогнозування за MSE або R^2 .

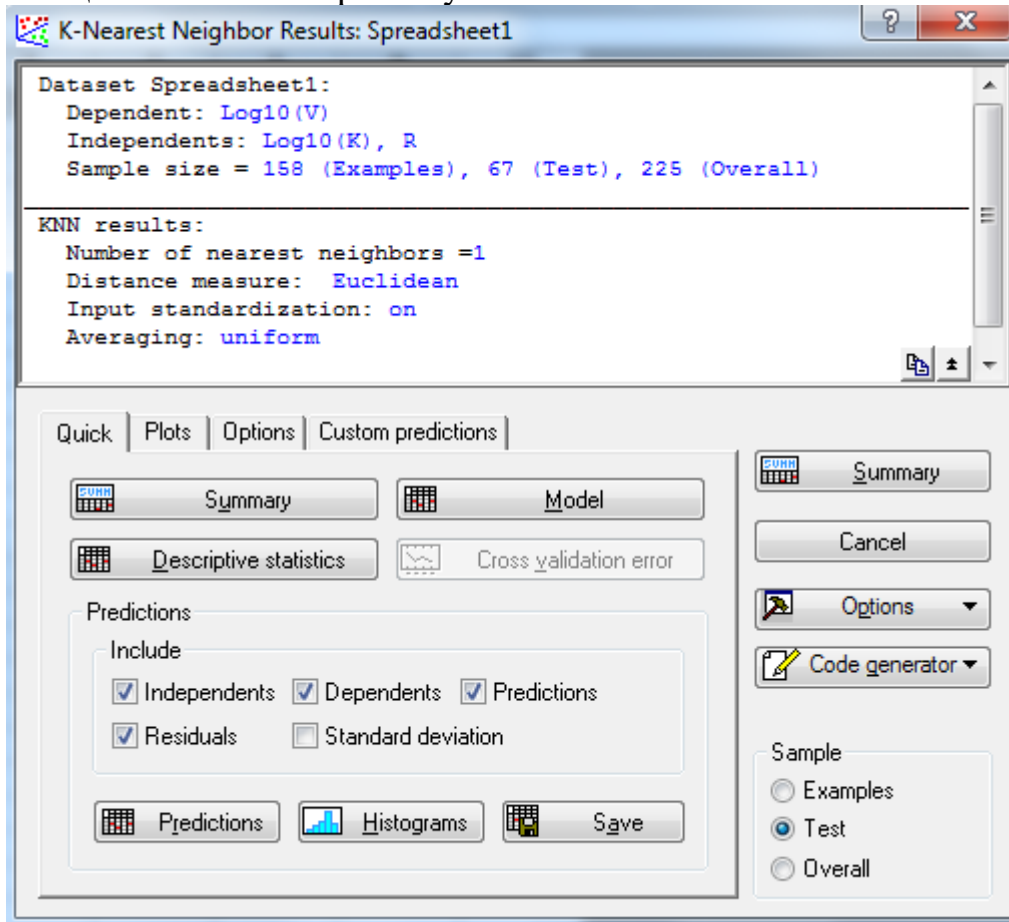


Рис. 2.3 – Результати прогнозування методом k -найближчих сусідів.

Крок 5. Порівняння моделей

1. Порівняйте похибки моделей SVM і k -NN.
2. Використайте таблицю для підсумку: Метод - MSE - R^2 - Переваги

Зміст звіту:

1. Тема, мета і завдання лабораторної роботи.
2. Короткі теоретичні відомості.
3. Етапи побудови моделей SVM та k -NN.
4. Графічні результати прогнозування.
5. Порівняння точності методів.
6. Висновки.

Контрольні запитання:

1. У чому полягає принцип роботи методу k -найближчих сусідів?
2. Як вибір параметра k впливає на точність моделі?

3. Які ядрові функції використовуються у SVM?
4. Які параметри впливають на узагальнюючу здатність SVM?
5. Які критерії використовуються для оцінки моделей у STATISTICA?
6. Як інтерпретувати графік Predicted vs Experimental?
7. У чому переваги й недоліки методів k -NN та SVM?

ЛАБОРАТОРНА РОБОТА №3

Тема: Застосування методу нейронних мереж до вирішення задач регресії з використанням навчального алгоритму зворотного поширення помилок у пакеті STATISTICA.

Мета роботи:

Закріпити теоретичні знання щодо побудови, навчання, оцінювання та застосування багатошарових нейронних мереж (MLP) для задач регресії; ознайомитися з інтерфейсом та можливостями пакету STATISTICA.

Короткі теоретичні відомості:

НМ - послідовність з'єднаних між собою нейронів, а нейрон - обчислювальна одиниця, яка отримує інформацію, виконує над нею прості математичні дії та передає її іншому нейрону. Кожен вхід нейрона, на який надходить деяка кількість сигналів, є виходом іншого. Кожен вхідний сигнал множать на відповідну вагу, аналогічну синаптичній силі, через що вхідна інформація змінюється під час передавання від одного нейрона до іншого. Далі всі результати додають і визначають рівень активації нейрона. НМ не запрограмовують, а навчають на даних.

Важливими параметрами НМ є її топологія, алгоритм навчання та функції активації прихованого та вихідного шарів. У цьому дослідженні застосовували алгоритм навчання Бройдена-Флетчера-Гольдфарба-Шанно (Broyden-Fletcher-Goldfarb-Shanno) (BFGS) та вибрали суму квадратів помилок (SOS) як функцію помилки. Для оцінювання точності моделі використовуються середньоквадратична похибка (MSE) та коефіцієнт детермінації (R^2).

Крок 1. Запуск системи STATISTICA

1. Відкрийте STATISTICA 12.
2. Оберіть інтерфейс Ribbon Bar або Classic Menu (рис. 3.1).
3. Для зміни інтерфейсу скористайтеся Options → Ribbon Bar або View → Ribbon Bar.

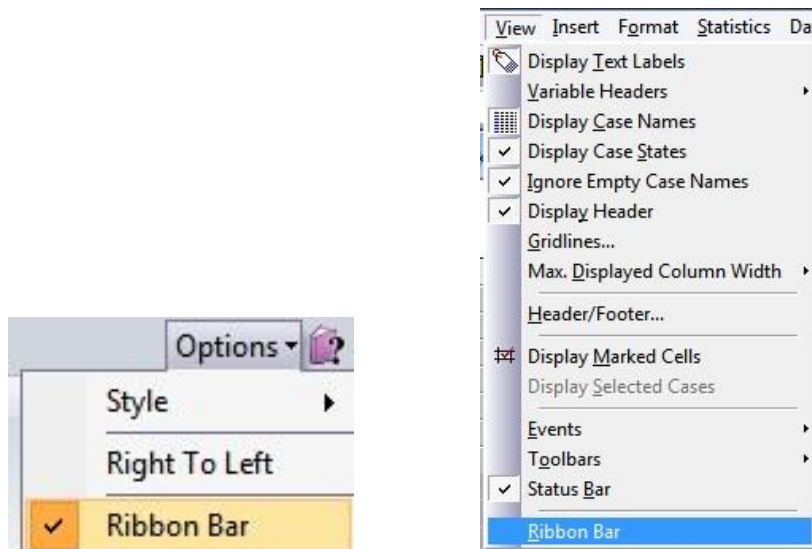


Рис. 3.1 – Вибір інтерфейсу STATISTICA.

Крок 2. Імпорт даних для регресійного аналізу

1. Виберіть File → Open → Data File.

2. Імпортуйте файл у форматі .STA, .XLS або .CSV.
3. Перевірте назви змінних: ΔK , R, da/dN тощо.
4. Збережіть таблицю у форматі STATISTICA Spreadsheet (.sta).



	1	2	3	4	5	6	7	8	9	10
	Log10(K)	Log10(V)	R	Var4	Var5	Var6	Var7	Var8	Var9	Var10
1	0,3579	-9,27	0,2							
2	0,361	-10,451	0,2							
3	0,3673	-9,7496	0,2							
4	0,3754	-9,3354	0,2							
5	0,3945	-9,0057	0,2							
6	0,3972	-9,0209	0,2							
7	0,4242	-8,9355	0,2							
8	0,4346	-8,983	0,2							
9	0,4396	-8,6635	0,2							
10	0,493	-8,4962	0,2							
11	0,4953	-8,6556	0,2							
12	0,5201	-8,6216	0,2							
13	0,5222	-8,4789	0,2							
14	0,5534	-8,4828	0,2							
15	0,5933	-8,3799	0,2							
16	0,5942	-8,251	0,2							
17	0,5959	-8,4123	0,2							
18	0,5977	-8,5171	0,2							
19	0,5986	-8,5229	0,2							
20	0,6012	-8,301	0,2							

Рис. 3.2 – Імпорт даних у STATISTICA.

Крок 3. Вибір типу аналізу

1. У головному меню виберіть Data Mining → Neural Networks → Regression (рис. 3.3).
2. У вікні Neural Networks натисніть Variables.
3. Вкажіть Dependent variable – змінну, яку потрібно передбачити; Covariates – незалежні змінні.
4. Натисніть ОК для підтвердження.

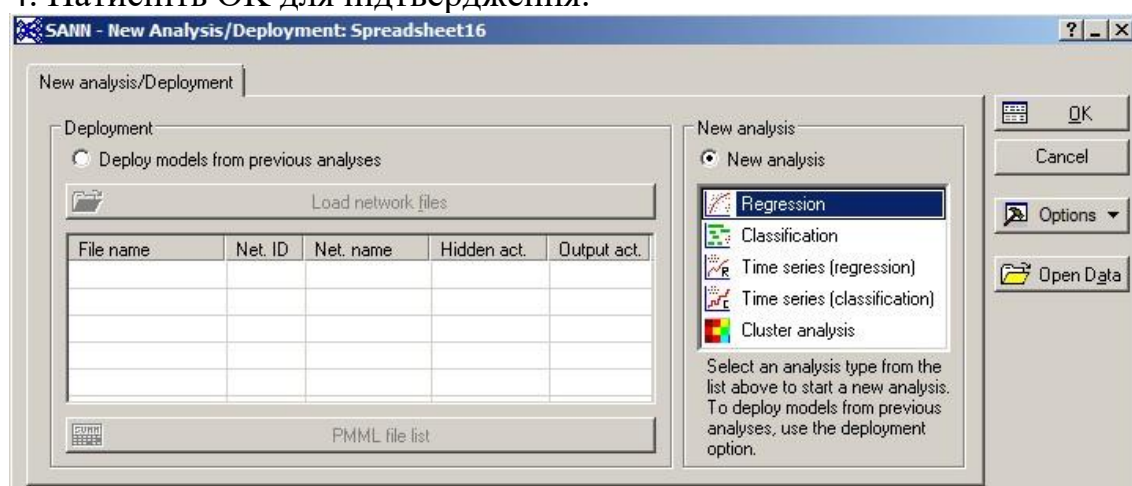


Рис. 3.3 – Вікно вибору змінних для задачі регресії.

Крок 4. Розподіл даних на навчальну та тестову вибірки

1. Виберіть вкладку Sampling.
2. Вкажіть пропорцію Training – 70%, Test – 30%.

3. Натисніть ОК для збереження параметрів.

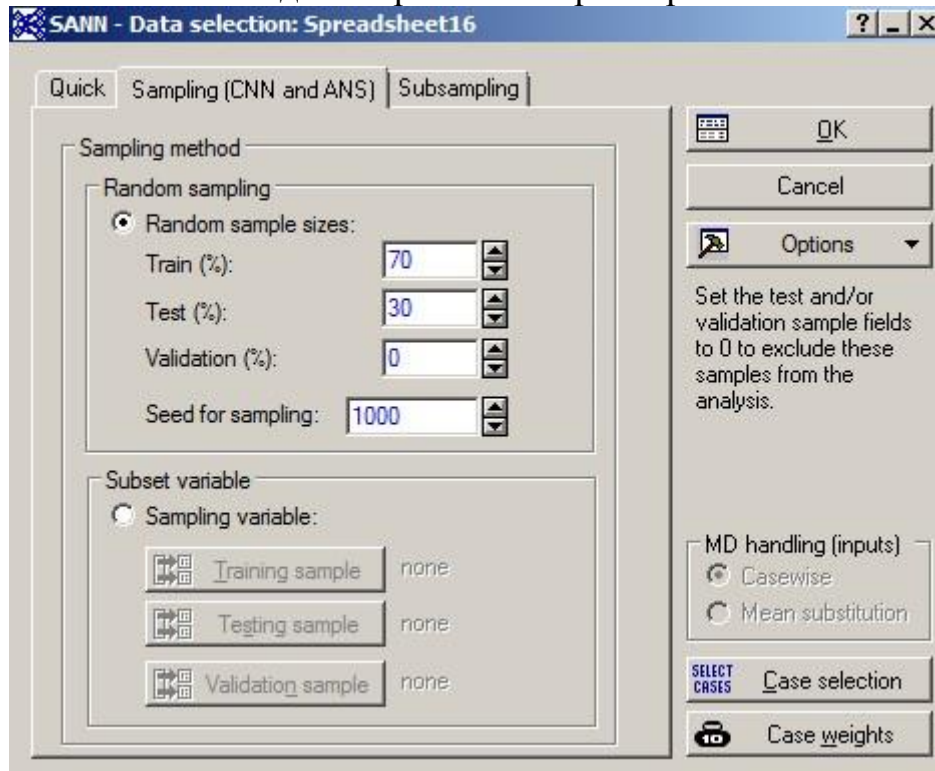


Рис. 3.4 – Вікно Sampling у STATISTICA.

Крок 5. Налаштування та навчання нейронної мережі

1. У меню Quick оберіть тип мережі MLP.
2. Вкажіть кількість нейронів у прихованому шарі (наприклад, 5–15).
3. Виберіть функції активації: Tanh або Sigmoid для прихованого шару, Linear для вихідного.
4. Натисніть Train, щоб запустити процес навчання.
5. Спостерігайте за зменшенням помилки під час навчання (рис. 3.5).

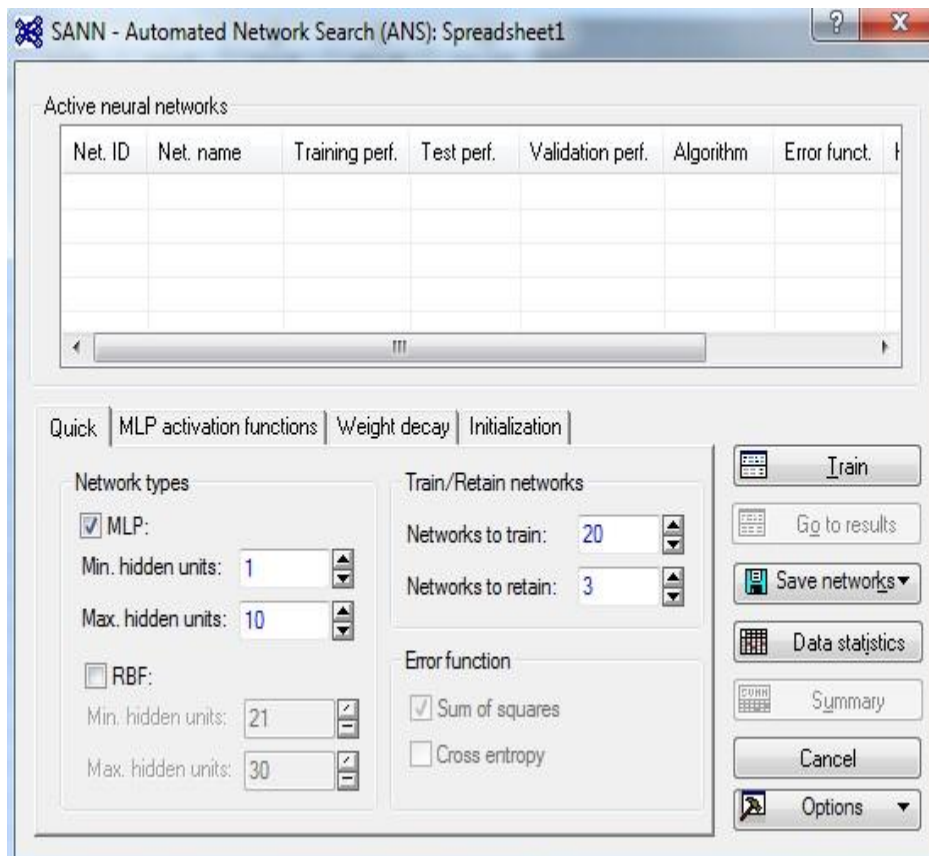


Рис. 3.5 – Вікно налаштування MLP у STATISTICA.

Крок 6. Аналіз результатів регресії

1. Перейдіть у вкладку Results → Predictions (рис. 3.6).
2. Перегляньте стовпці Experimental, Predicted, Residual.
3. Побудуйте графік Predicted vs Experimental для оцінки якості моделі.
4. Обчисліть похибку MSE.

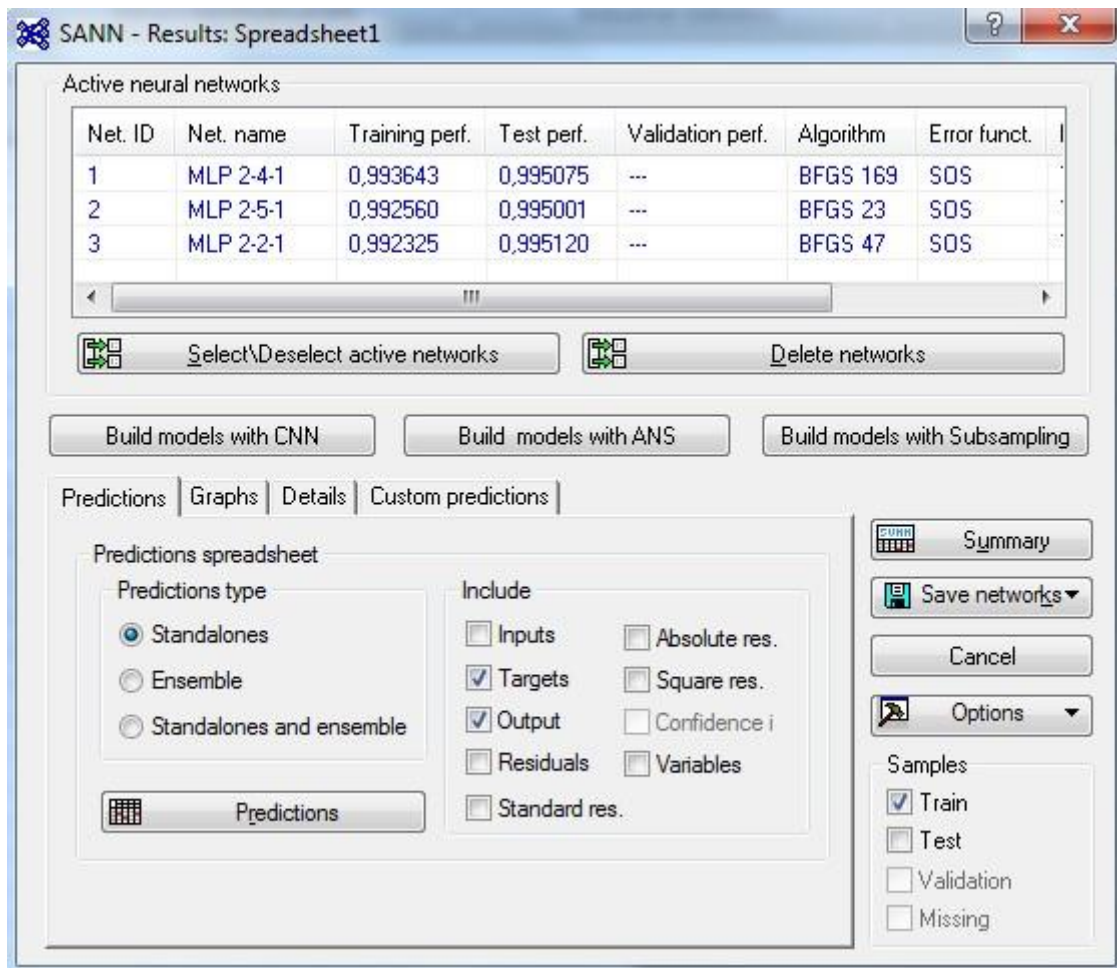


Рис. 3.6 – Результати прогнозування методом MLP.

Крок 7. Підсумковий аналіз і збереження моделі

1. Збережіть навчену мережу через File → Save Network As.
2. Порівняйте результати декількох моделей за показниками MSE та R^2 .
3. Зробіть висновок щодо ефективності моделі.

Зміст звіту:

1. Тема, мета і завдання лабораторної роботи.
2. Короткі теоретичні відомості.
3. Етапи побудови регресійної моделі у STATISTICA.
4. Результати навчання та прогнозування.
5. Графік Predicted vs Experimental.
6. Обчислення MSE, R^2 .
7. Висновки.

Контрольні запитання:

1. Що таке багатошарова нейронна мережа (MLP)?
2. Які функції активації використовуються для регресійних задач?
3. У чому полягає алгоритм зворотного поширення помилки?
4. Як здійснюється розподіл даних на навчальну та тестову вибірки?
5. Що таке середньоквадратична похибка (MSE)?
6. Як інтерпретувати коефіцієнт детермінації (R^2)?

7. Як інтерпретувати графік Predicted vs Experimental?
8. Як уникнути перенавчання при побудові регресійних моделей?

ЛАБОРАТОРНА РОБОТА №4

Тема: Застосування методу нейронних мереж до вирішення задач класифікації з використанням навчального алгоритму зворотного поширення помилок у пакеті STATISTICA.

Мета роботи:

Закріпити теоретичні знання щодо побудови, навчання, оцінювання та застосування багатошарових нейронних мереж (MLP) для задач класифікації; ознайомитися з інтерфейсом та можливостями пакету STATISTICA.

Короткі теоретичні відомості:

Багатошарова нейронна мережа (MLP) складається з вхідного, прихованого та вихідного шарів. Для задач класифікації вихідний шар зазвичай використовує сигмоїдну або softmax функцію активації. Алгоритм зворотного поширення помилок (Back Propagation) мінімізує різницю між прогнозованими та реальними класами, оновлюючи ваги зв'язків між нейронами. Розділення даних на навчальну та тестову вибірки дозволяє перевірити узагальнюючу здатність моделі.

Крок 1. Запуск системи STATISTICA

1. Відкрийте STATISTICA 12.
2. Оберіть інтерфейс Ribbon Bar або Classic Menu (рис. 4.1).
3. Для зміни інтерфейсу скористайтеся Options → Ribbon Bar або View → Ribbon Bar.

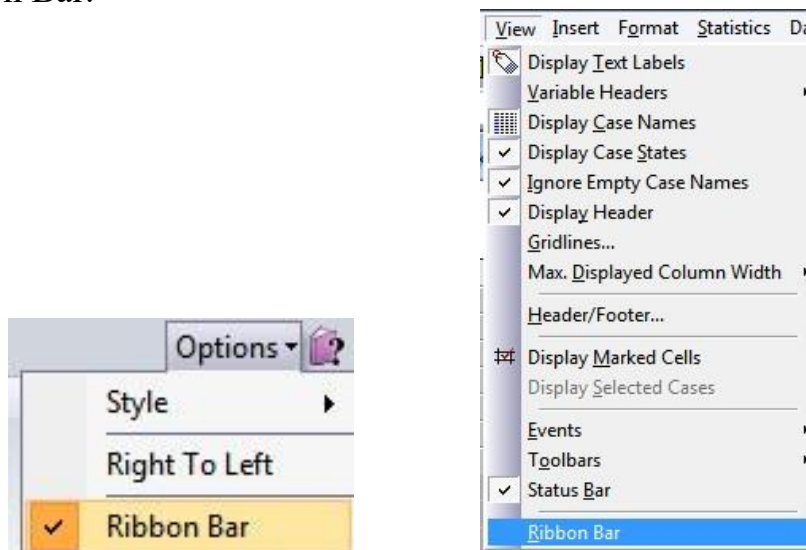


Рис. 4.1 – Вибір інтерфейсу STATISTICA.

Крок 2. Імпорт даних для класифікації

1. Виберіть File → Open → Data File.
2. Імпортуйте файл у форматі .STA, .XLS або .CSV.
3. Перевірте коректність назв змінних і типів даних.
4. Збережіть таблицю у форматі STATISTICA Spreadsheet (.sta).

Крок 3. Вибір типу аналізу

1. У головному меню виберіть Data Mining → Neural Networks → Classification (рис. 4.2-4.3).

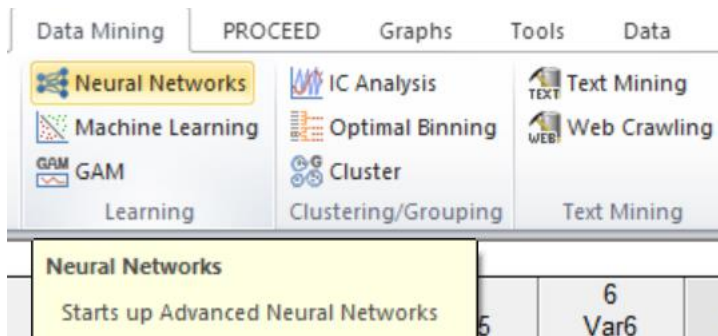


Рис. 4.2 – Вікно Data Mining

2. У вікні Neural Networks натисніть Variables.
3. Задайте Dependent variable – клас, що прогнозується; Covariates – вхідні змінні.
4. Натисніть ОК для підтвердження.

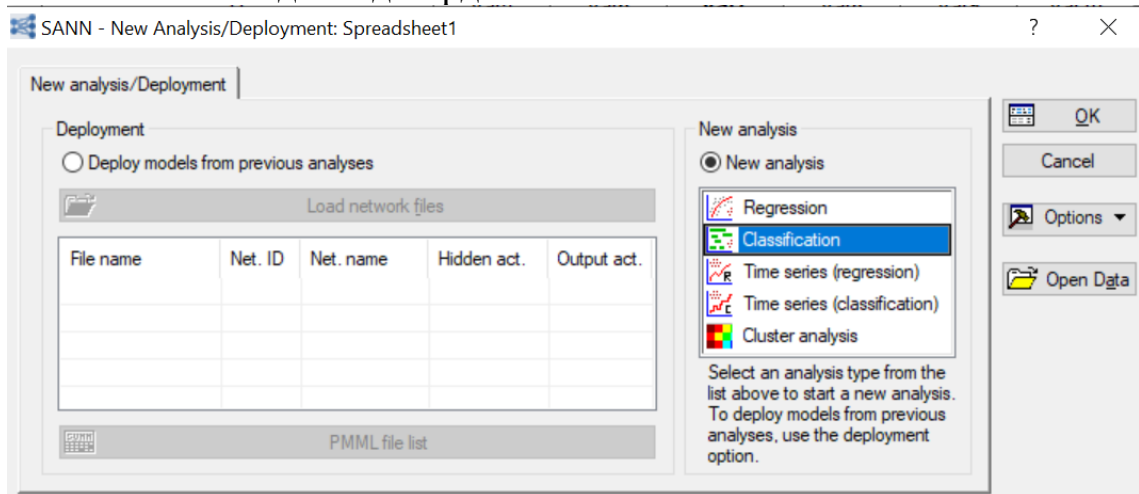


Рис. 4.3 – Вікно вибору змінних для задачі класифікації.

Крок 4. Розподіл даних на навчальну та тестову вибірки

1. Виберіть вкладку Sampling.
2. Вкажіть пропорцію Training – 70%, Test – 30%.
3. Натисніть ОК для збереження параметрів.

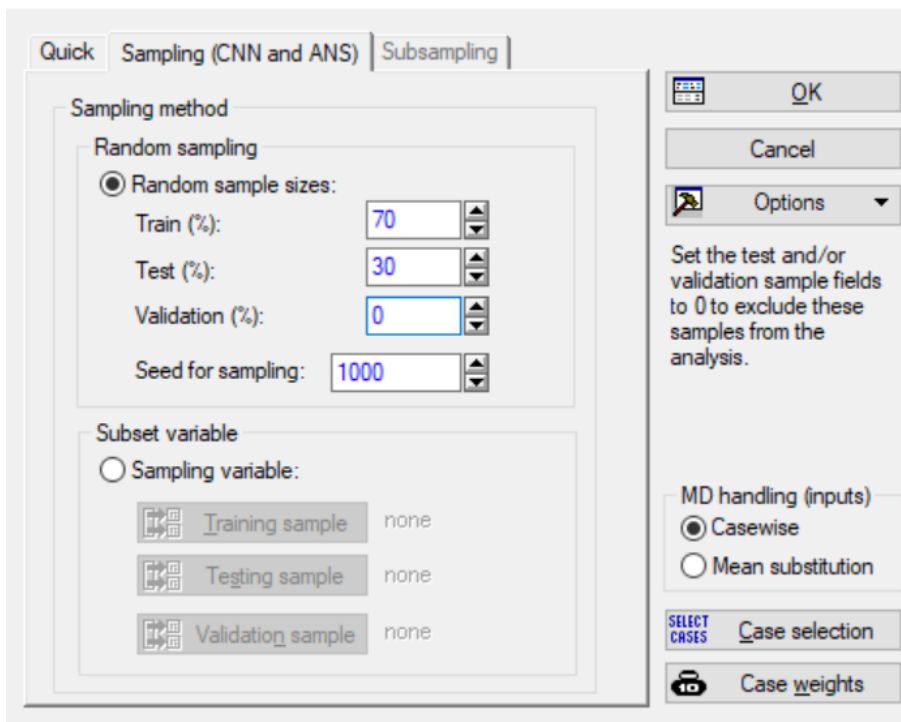


Рис. 4.4 – Вікно Sampling у STATISTICA.

Крок 5. Налаштування та навчання нейронної мережі

1. У меню Quick оберіть тип мережі MLP.
2. Вкажіть кількість нейронів у прихованому шарі (наприклад, 5–30).
3. Виберіть функції активації: Sigmoid, наприклад, для прихованого шару, Softmax для вихідного.
4. Натисніть Train, щоб запустити процес навчання.
5. Спостерігайте за зміною помилки під час навчання (рис.4.5).

Neural network training in progress...

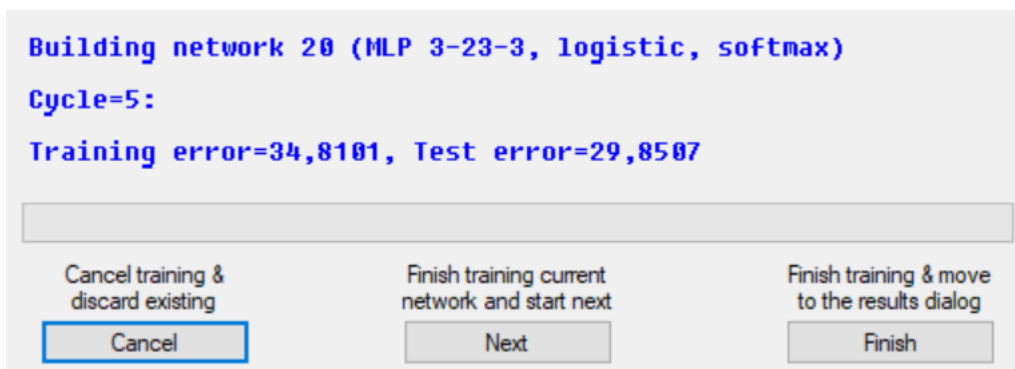


Рис. 4.5 – Вікно налаштування MLP у STATISTICA.

Крок 6. Аналіз результатів класифікації

1. Перейдіть у вкладку Results → Classification Results (рис. 4.6).
2. Перегляньте таблицю Confusion Matrix.
3. Обчисліть точність класифікації: Accuracy, Precision, Recall, F1 score, тощо.
4. Перегляньте графік ROC Curve для оцінки якості класифікації.

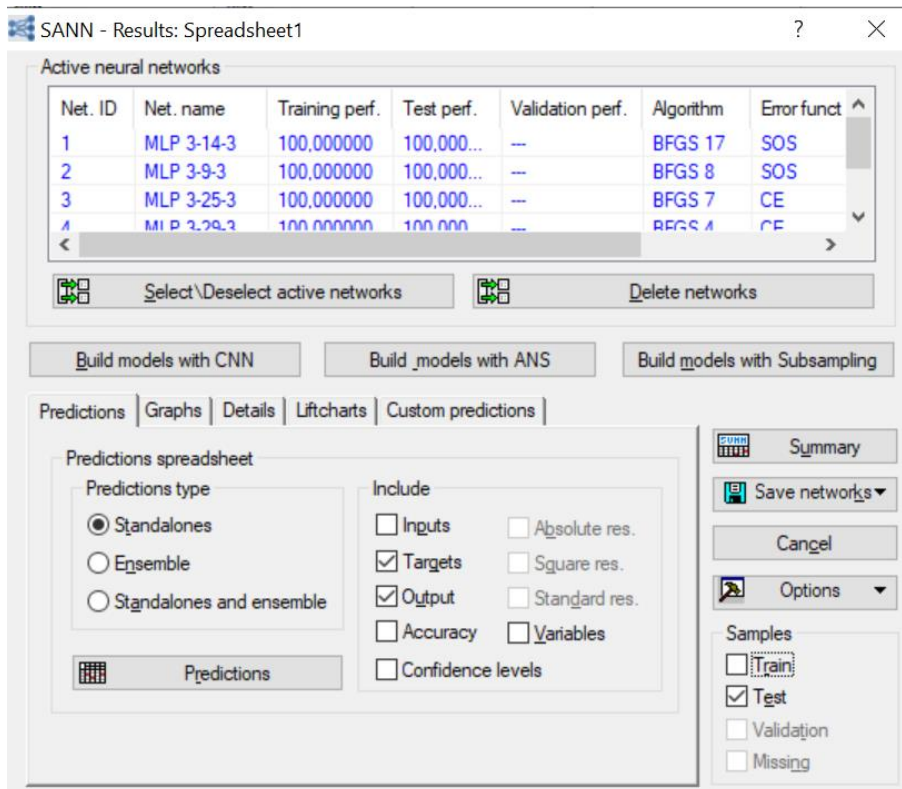


Рис. 4.6 – Вікно результатів класифікації у STATISTICA.

Крок 7. Підсумковий аналіз і збереження моделі

1. Збережіть навчену мережу через File → Save Network As.
2. Порівняйте результати декількох моделей за показником точності.
3. Зробіть висновок щодо ефективності навчання.

Зміст звіту:

1. Тема, мета і завдання лабораторної роботи.
2. Короткі теоретичні відомості.
3. Етапи побудови моделі класифікації у STATISTICA.
4. Результати навчання та класифікації.
5. Confusion Matrix, ROC-графіки.
6. Висновки.

Контрольні запитання

1. Що таке багатошарова нейронна мережа (MLP)?
2. Як працює алгоритм зворотного поширення помилки?
3. Які функції активації використовуються для класифікації?
4. Як розподіляються дані на навчальну та тестову вибірки?
5. Що таке Confusion Matrix і як вона інтерпретується?
6. Які показники використовуються для оцінки точності класифікації?
7. Що показує ROC-крива?
8. Як уникнути перенавчання нейронної мережі?

ЛАБОРАТОРНА РОБОТА №5

Тема: Генерація дизайну пляшкових корків за допомогою дифузійних моделей Stable Diffusion та LoRA

Мета роботи: Ознайомитися з принципами роботи дифузійних моделей генерації зображень, навчитися розгортати середовище Google Colab для практичного застосування Stable Diffusion із LoRA-донавачанням та реалізувати власний Gradio-інтерфейс для генерації дизайнів пляшкових корків.

Теоретичні відомості:

Дифузійні моделі - це клас генеративних нейронних мереж, що відтворюють дані шляхом зворотного процесу дифузії шуму. Вони складаються з прямого процесу (додавання шуму до зображення) та зворотного процесу (відновлення зображення). Stable Diffusion - латентна дифузійна модель, яка працює в стисненому просторі ознак. LoRA - метод адаптації великих моделей на малих вибірках без повного перенавчання. Він змінює лише частину ваг у заданих шарах, що зменшує обчислювальні витрати, зберігає стабільність базової моделі й дозволяє швидко додавати нові стилі генерації.

Сучасне виробництво пляшкових ковпачків, зокрема термоусадочних і алюмінієвих, передбачає реалізацію високотехнологічних процесів, що забезпечують точність, повторюваність та якість продукції відповідно до міжнародних стандартів. У межах дослідження було проаналізовано виробничий процес на підприємстві АТ «Технологія» - провідному українському виробнику пакувальної продукції. Використання поліамінатних ковпачків дозволяє значно розширити дизайнерські можливості завдяки металізованій поверхні, яка забезпечує глибокі, насичені кольори та створює ефект преміальності. Такі властивості вимагають точного візуального відтворення при генерації зображень, що підвищує вимоги до моделі штучного інтелекту, зокрема, вона має не лише відображати загальні риси, але й зберігати текстуру, блиск і структурну автентичність матеріалу.

Для моделювання процесу створення дизайну пляшкових корків в умовах обмеженої кількості прикладів було застосовано генеративний підхід, що відтворює візуальну складність, багатошаровість оформлення й варіативність стилів. Ефективне навчання нейронної мережі для генерації дизайнів пляшкових корків базується на ретельно сформованій навчальній вибірці. До навчальної вибірки було включено 23 зразки пляшкових корків підприємства АТ «Технологія», котрі відображають різні типи ковпачків. Кожне зображення мало високу роздільну здатність та стандартизоване тло (однорідне світле або біле), що сприяло кращій генералізації під час тренування. Додатково, кожному зразку було надано опис, що узагальнює його ключові візуальні характеристики, а саме, кольори, форми, текстури, наявність логотипів чи написів. Наприклад, "black and gold cap with floral band and 'Monte Choco' text" (для зразка 4) чи "shiny gold bottle cap with 'VP' logo on the top" (для зразка 17). Такі анотації використовувались як текстові підказки під час генерації нових зображень і відігравали критичну роль у контекстному навчанні моделей, що

реалізують механізми перехресної уваги (cross-attention) для узгодження тексту та зображення.

Крок 1. Підготовка середовища Google Colab

1. Відкрити сторінку <https://colab.research.google.com>.
2. Створити новий ноутбук і обрати GPU як тип обчислювального середовища (*Runtime* → *Change runtime type* → *GPU*).
3. У першій комірці встановити необхідні бібліотеки (Diffusers, Transformers, Accelerate, Safetensors, Gradio).
4. Перевірити, що GPU підключено успішно.

Крок 2. Завантаження базової моделі Stable Diffusion

1. Завантажити модель **Stable Diffusion v1.5** із відкритого джерела.
2. Ініціалізувати модель у середовищі Colab та перевірити її готовність до роботи.
3. Після завантаження модель готова до подальшого донавчання або генерації.

Крок 3. Підключення LoRA-ваг

1. Підготувати або завантажити попередньо натреновані ваги LoRA.
2. Підключити LoRA до Stable Diffusion через інтерфейс Google Colab.
3. Переконавшись, що LoRA активована - це дозволяє моделі використовувати стиль вашого навчального набору.

Крок 4. Генерація зображень

1. Задати текстовий опис бажаного дизайну пляшкового корка (prompt).
2. Виконати генерацію, використовуючи параметри:
3. Отримати перші результати генерації і зберегти найкращі зображення у форматі PNG.



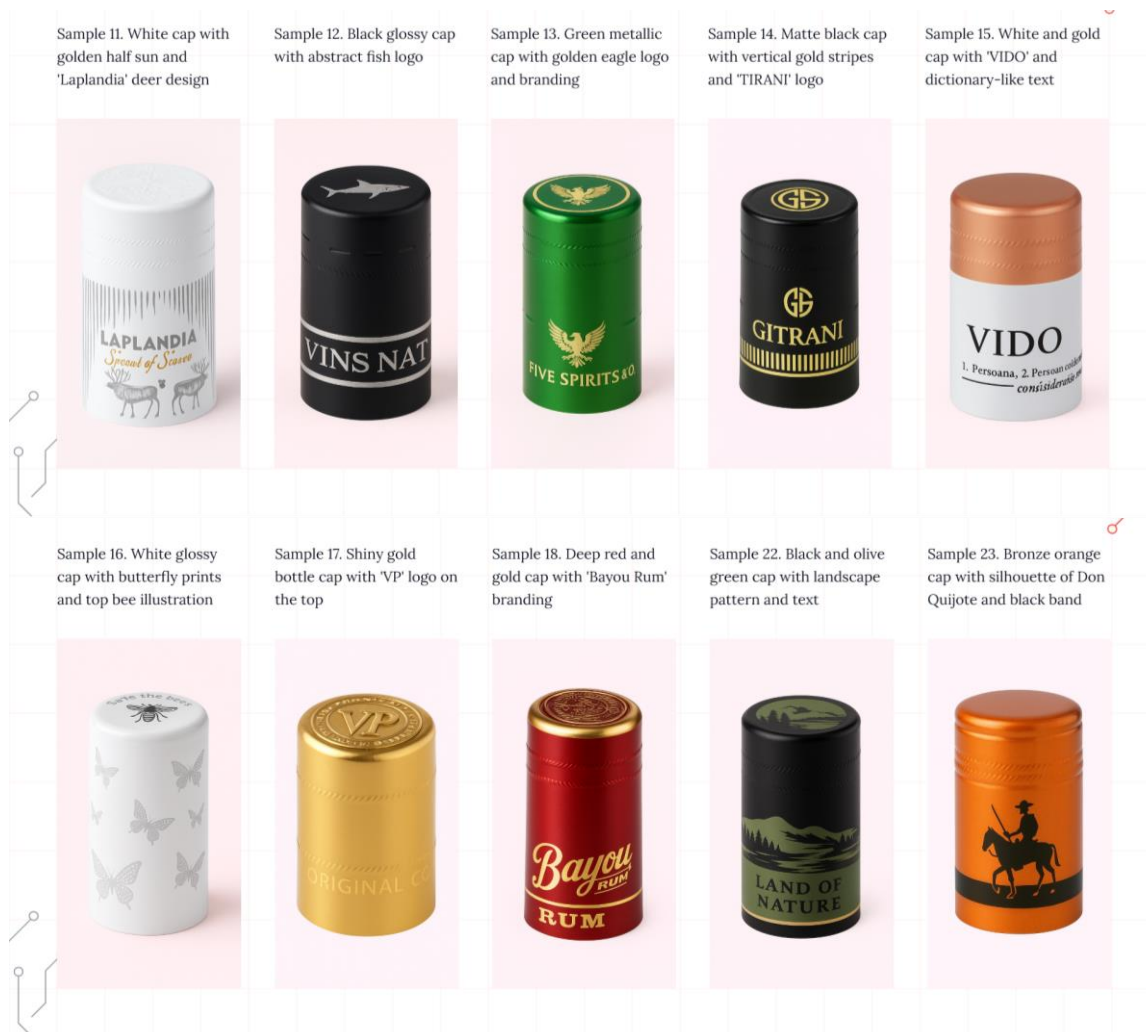


Рис. 5.1 – Приклад згенерованого дизайну пляшкового корка.

Крок 5. Створення веб-інтерфейсу Gradio

1. Створити просту форму введення для опису (prompt) та зображення результату.
2. Додати назву інтерфейсу, наприклад: *Bottle Cap Design Generator*.
3. Запустити Gradio-додаток у середовищі Google Colab та протестувати його роботу.
4. Користувач може змінювати опис і отримувати різні варіанти дизайну.

Enter your prompt:

black and gold cap with floral band and
'Monte Choco' text

Submit

Generated image:



Рис. 5.2 – Веб-інтерфейс Gradio для генерації дизайну пляшкових корків.

Крок 6. Аналіз результатів генерації

1. Змінювати параметри генерації та спостерігати, як вони впливають на вигляд і деталізацію зображень.
2. Визначити найоптимальніші параметри для отримання якісного результату.
3. Зробити порівняння кількох варіантів дизайну.
4. Відібрати найкращі результати для звіту.

Крок 7. Збереження результатів та підсумковий аналіз

1. Завантажити готові результати через меню **Files** → **Download**.
2. Зберегти ноутбук Google Colab для подальшого використання або відновлення моделі.
3. Оцінити ефективність генерації та зробити висновки щодо можливостей LoRA та Stable Diffusion у дизайні.

Зміст звіту

1. Тема, мета та завдання лабораторної роботи.
2. Короткі теоретичні відомості (Stable Diffusion, LoRA, Gradio).
3. Покроковий опис виконання в середовищі Google Colab.
4. Приклади згенерованих зображень.
5. Пояснення впливу параметрів генерації на результат.
6. Висновки.

Контрольні запитання:

1. Що таке дифузійна модель і який принцип її роботи?

2. Які переваги має модель Stable Diffusion?
3. У чому полягає метод LoRA та які його переваги при донавчанні моделей?
4. Які параметри впливають на якість генерації зображень?
5. Яку роль відіграють параметри *guidance scale* та *num inference steps*?
6. Для чого використовується Gradio-інтерфейс?
7. Чому Google Colab є зручним середовищем для генеративного штучного інтелекту?
8. Які можливості відкриває використання LoRA у сфері автоматизованого дизайну?

ЛАБОРАТОРНА РОБОТА №6

Тема: Застосування нейромережових платформ для формування зображень на основі тексту (*DALL·E 3, Midjourney, Leonardo, Firefly, ImageFX*)

Мета роботи: Ознайомитися з принципами роботи сучасних генеративних нейронних мереж, навчитися формувати якісні текстові запити (prompts) та порівняти результати роботи популярних платформ штучного інтелекту для генерації зображень.

Короткі теоретичні відомості:

Генеративні моделі зображень - це системи штучного інтелекту, здатні створювати нові візуальні зразки на основі текстового опису. Найпоширеніші технології, що використовуються сьогодні:

- DALL·E 3 (OpenAI) - інтегрована в ChatGPT модель для створення високоякісних реалістичних зображень.
- Midjourney - працює через Discord, дозволяє генерувати художні та стильові композиції з великою деталізацією.
- Leonardo AI - орієнтований на дизайнерів, підтримує створення власних LoRA-моделей і комбінування стилів.
- Adobe Firefly - сфокусований на творчих зображеннях і текстових ефектах, тісно інтегрований з Photoshop.
- ImageFX (Google) - новітня платформа, що використовує дифузійні алгоритми для генерації з високою точністю деталей.

Всі ці сервіси реалізують підхід text-to-image, у якому текстові описи перетворюються на візуальні образи через глибокі нейронні архітектури.

Крок 1. Вибір платформи для роботи

1. Ознайомитися з можливостями генеративних сервісів: DALL·E 3, Midjourney, Leonardo, Firefly, ImageFX.
2. Обрати платформи для виконання лабораторної роботи.
3. Зареєструватися або увійти до обраної системи.

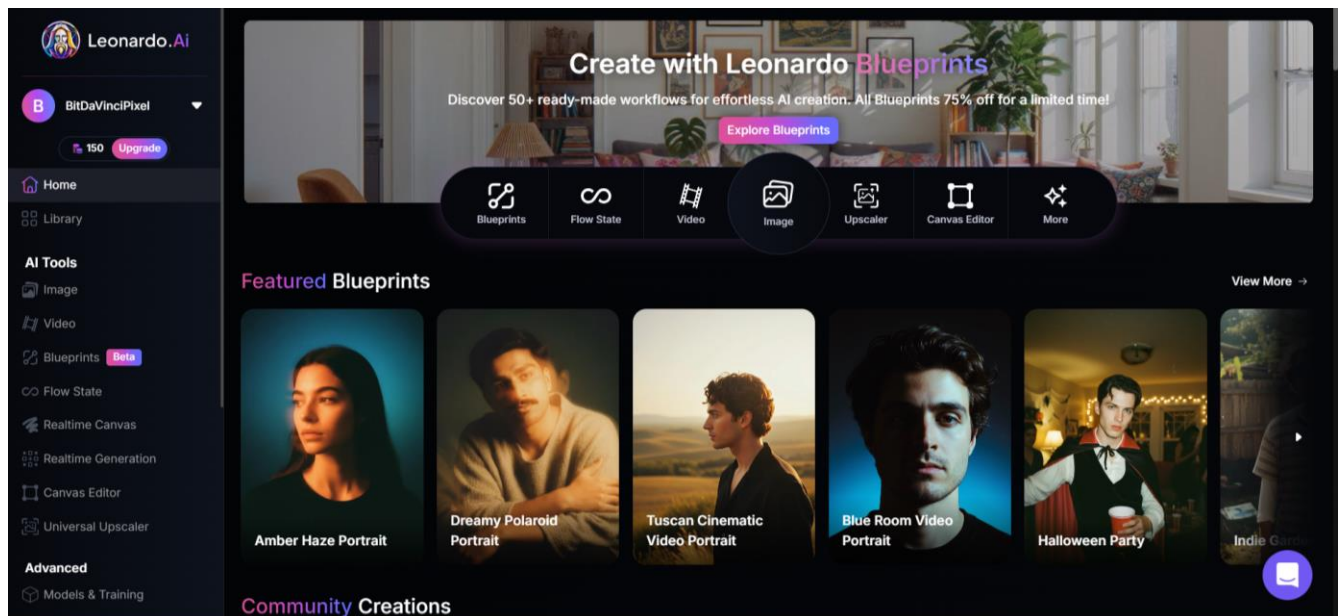


Рис. 6.1 – Інтерфейси сучасних генеративних платформ.

Крок 2. Формулювання текстового запиту (prompt)

1. Створити короткий опис бажаного зображення.
Наприклад: *«Створити високоякісне, реалістичне зображення Лесі Українки в традиційній українській вишиванці, відображаючи сучасний дух і культурну ідентичність. Вона повинна стояти на оживленому місці в сучасній Україні, на фоні Хрещатика. На задньому плані повинен бути видимий напис “Україна”, вбудований у міський пейзаж, підкреслюючи зв'язок між фігурою та місцем. Кожна деталь повинна бути ретельно відтворена.»*
2. Додати уточнення щодо стилю (реалістичний, анімаційний, 3D), кольорів і матеріалів.
3. Використовувати англійську мову для більш точного результату.

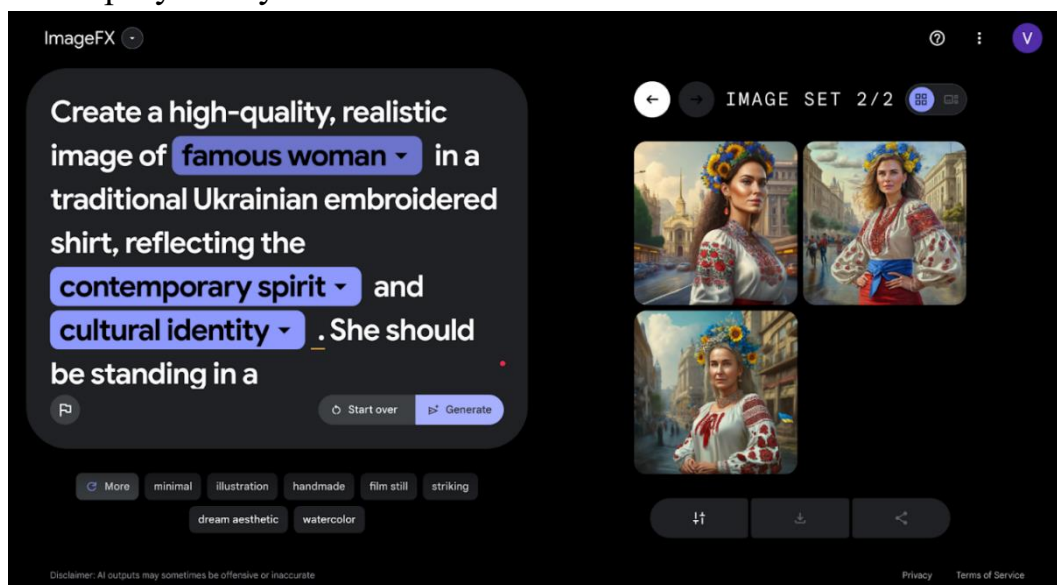


Рис. 6.2 – Приклад створення текстового запиту.

Крок 3. Генерація зображень

1. Увести опис у поле prompt на вибраній платформі.
2. Запустити процес генерації.
3. Дочекатися відображення результатів і вибрати найкращі варіанти.



Рис. 6.3 – Результат генерації в DALL·E 3.



Рис. 6.4 – Результат генерації в Midjourney.

Крок 4. Порівняння результатів між платформами

1. Згенерувати однакові запити в різних системах.
2. Порівняти якість, освітлення, колір, деталізацію та художній стиль.

3. Визначити, яка платформа найточніше відтворює задум.

Крок 5. Оптимізація prompt-запитів

1. Вдосконалити опис, додаючи більше деталей (освітлення, композицію, стиль).
2. Проаналізувати, як зміни prompt впливають на зображення.
3. Визначити найефективнішу структуру запити для досягнення бажаного результату.

Крок 6. Аналіз отриманих результатів

1. Зберегти всі варіанти для подальшого аналізу.
2. Визначити, які типи сцен краще генеруються кожною платформою.
3. Відмітити сильні та слабкі сторони кожної системи.

Крок 7. Висновки та рекомендації

1. Сформулювати підсумкові висновки щодо ефективності генеративних платформ.
2. Вказати, які інструменти виявилися найбільш зручними.
3. Надати поради для створення якісних prompt-запитів.

	I mageFX	A dobe Firefly	L eonardo	Midj ourney	DA LL-E 3
Representation of Lesya Ukrainka's face	+	+	+	+	+
Embroidery details	+	+	+	+	+
Clarity of the background	-	-	-	+	+
Generation of text in the background	-	-	-	-	+

Рис. 6.5 – Підсумкове порівняння платформ за критеріями якості.

Зміст звіту:

1. Тема, мета і завдання лабораторної роботи.
2. Короткі теоретичні відомості.
3. Покроковий опис виконання роботи.
4. Приклади отриманих результатів.
5. Порівняльний аналіз платформ.
6. Висновки.

Контрольні запитання:

1. Які існують популярні платформи для генерації зображень на основі тексту?
2. У чому полягає принцип роботи дифузійних моделей?

3. Які особливості має модель DALL·E 3?
4. Як правильно формулювати prompt-запит?
5. Чим відрізняються Midjourney і Leonardo AI за підходом до генерації?
6. Які критерії можна використати для оцінки якості зображень?
7. Як можна покращити результат генерації?
8. У яких галузях можна застосовувати такі платформи?

РЕКОМЕНДОВАНА ЛІТЕРАТУРА

1. Штучний інтелект. Нейромережева обробка інформації : архітектури, навчання, застосування : навчальний посібник у 2-х ч. : Ч. 1 / О. Г. Руденко, О. О. Безсонов, С. П. Євсєєв, О. Б. Ахієзер, Ю. І. Зайцев ; за заг. ред. С. П. Євсєєва. – Харків : НТУ «ХПІ», – Львів : «Новий Світ-2000», 2025. – 426 с. – (Серія «Кібербезпека та штучний інтелект»). ISBN 978-966-418-517-9
2. Штучний інтелект. Нейромережева обробка інформації : архітектури, навчання, застосування : навчальний посібник у 2-х ч. : Ч. 2 / О. Г. Руденко, О. О. Безсонов, С. П. Євсєєв, О. Б. Ахієзер, Ю. І. Зайцев ; за заг. ред. С. П. Євсєєва. – Харків : НТУ «ХПІ», – Львів : «Новий Світ-2000», 2025. – 376 с. – (Серія «Кібербезпека та штучний інтелект»). ISBN 978-966-418-518-6
3. Глибовець М.М., Олецький О.В. Штучний інтелект. Підручник для студентів вищих навчальних закладів, які навчаються за спеціальностями "Комп'ютерні науки" та "Прикладна математика". - К.:Вид.дім "КМ Академія", 2002. - 366 с.
4. Глибовець М.М., Кравченко І.В., Олецький О.В., Терещенко В.М. Програмування в Пролозі. Навчальний посібник. - К.: ВЦ "Київський університет", 1998. - 110 с.
5. Гладун А. Я., Рагушина Ю. В. Data Mining: пошук знань в даних. К.: ТОВ «ВД «АДЕФ Україна», 2016. 452 с.
6. Литвин В. В., Пасічник В. В., Нікольський Ю. В. Аналіз даних та знань : навчальний посібник. Львів: «Магнолія 2006», 2015. 276 с.
7. Снитюк В. Є. Прогнозування. Моделі. Методи. Алгоритми : навчальний посібник. К.: Маклаут, 2008. 364 с.
8. Литвин В.В., Пасічник В.В., Яцишин Ю.В. Інтелектуальні системи : підручник. Львів: Новий світ – 2000, 2009. 406с.
9. Математичне забезпечення інформаційно-керуючих систем: підручник / А. М. Гуржій, З. В. Дудар, В. М. Левикін, Б. В. Шамша. Х. : Компанія Сміт, 2006. 448с.
10. Haykin, S., 2009 “Neural Networks and Learning Machines”, Third Edition, McMaster University, Hamilton, Ontario, Canada, 938
11. Alpayndin, E. (2010). Introduction to machine learning. The Knowledge Engineering Review, 25(3), 353–353.
12. Yasniy, O., Mytnyk, M., Maruschak, P., Mykityshyn, A., & Didych, I. (2024). Machine learning methods as applied to modelling thermal conductivity of

- epoxy-based composites with different fillers for aircraft. *Aviation*, 28(2), 64-71.
13. Yasniy, O., Menou, A., Mykytyshyn, A., Kubashok, V., Didych, I. (2024). Application of neural network platforms for text-based image generation. *CEUR Workshop Proceedings*, 3842, pp. 232–240.
 14. Tymoshchuk, D., Yasniy, O., Maruschak, P., Iasnii, V., & Didych, I. (2024). Loading Frequency Classification in Shape Memory Alloys: A Machine Learning Approach. *Computers*, 13(12), 339.
 15. Iryna, D., Andrii, M., Andrii, S., & Mykola, M. (2024). Application of machine learning methods to the prediction of NO₂ concentration in the air environment. In *Ceur Workshop Proceedings Open source preview*, 2024, 3896, pp. 569–577.
 16. Iryna Didych, Oleh Yasniy. (2025). Structural integrity & Lifetime estimation by machine learning methods. Important structural elements. LAP Lambert Academic Publishing, Republik of Moldova, Europe. 124 p.
 17. Yasniy, O., Maruschak, P., Mykytyshyn, A., Didych, I., & Tymoshchuk, D. (2025). Artificial intelligence as applied to classifying epoxy composites for aircraft. *Aviation*, 29(1), 22-29.
 18. Tymoshchuk, D., Didych, I., Maruschak, P., Yasniy, O., Mykytyshyn, A., & Mytnyk, M. (2025). Machine Learning Approaches for Classification of Composite Materials. *Modelling*, 6(4), 118.

[illegible]

