

Дипломна робота

Деякі застосування байєсівських методів у машинному навчанні

Виконав: Бондаренко Є.С., КА-21

Науковий керівник: к. ф.-м. н., доц. Каніовська І. Ю.

Дипломна робота

- **Об'єкт дослідження:** алгоритми машинного навчання.
- **Предмет дослідження:** байєсівський статистичний підхід та його застосування в задачах машинного навчання.
- **Мета роботи:** проаналізувати предмет дослідження та дослідити можливість покращення результативності деяких алгоритмів машинного навчання за допомогою байєсівських методів.
- **Метод дослідження:** методи штучного інтелекту, теорії ймовірностей, статистичного аналізу, лінійної алгебри, дискретної математики, оптимізації (чисельні та аналітичні).

Актуальність дослідження

- Сьогодні, у добу big data, алгоритми **машинного навчання** є надзвичайно популярними:
 - біоінформатика
 - обробка природньої мови
 - комп'ютерний зір
 - пошукові двигуни
 - розпізнавання образів
 - рекомендаційні системи
 - Такі ІТ-гіганти як **Google, Microsoft, Facebook** та багато інших вкладають щорічно багато коштів у дослідження в області машинного навчання.
 - У свою чергу, **байєсівські методи** дозволяють вирішити велику кількість проблем та покращити результативність алгоритмів машинного навчання.
- 
- A word cloud in the top right corner of the slide. The central and largest text is 'MACHINE LEARNING'. Surrounding it are various related terms in different sizes and orientations, including 'SUPERVISED', 'UNSUPERVISED', 'NEURAL NETWORK', 'DATA SCIENCE', 'ARTIFICIAL INTELLIGENCE', 'STATISTICS', 'MODELS', 'ALGORITHMS', 'FEATURES', 'LABELS', 'CLASSES', 'CLUSTERS', 'DECISION TREES', 'RANDOM FORESTS', 'SUPPORT VECTOR MACHINES', 'LOGIT', 'GLM', 'BAYESIAN', 'PROBABILISTIC', 'DISTRIBUTION', 'METHODS', 'PERFORMANCE', 'EVALUATION', 'METRICS', 'CONFUSION MATRIX', 'PRECISION', 'RECALL', 'F1 SCORE', 'AUC', 'ROC CURVE', 'CROSS-VALIDATION', 'OVERFITTING', 'UNDERFITTING', 'BIAS-VARIANCE TRADEOFF', 'REGULARIZATION', 'LASSO', 'RIDGE', 'Elastic Net', 'Gradient Descent', 'Stochastic Gradient Descent', 'Adam', 'AdaGrad', 'RMSprop', 'Nesterov Momentum', 'Batch Normalization', 'Dropout', 'Data Augmentation', 'Transfer Learning', 'Domain Adaptation', 'Federated Learning', 'Reinforcement Learning', 'Deep Reinforcement Learning', 'Monte Carlo Tree Search', 'AlphaGo', 'AlphaZero', 'Self-Play', 'Evolutionary Algorithms', 'Genetic Algorithms', 'Simulated Annealing', 'Tabular Data', 'Text Data', 'Image Data', 'Audio Data', 'Sensor Data', 'Time Series Data', 'Network Data', 'Social Network Data', 'Geospatial Data', 'Healthcare Data', 'Financial Data', 'Marketing Data', 'Operational Data', 'Log Data', 'Clickstream Data', 'Web Analytics Data', 'Mobile App Data', 'IoT Data', 'Big Data', 'Cloud Data', 'Data Lake', 'Data Warehouse', 'Data Mart', 'Data Hub', 'Data Catalog', 'Data Governance', 'Data Privacy', 'Data Security', 'Data Quality', 'Data Integration', 'Data Interoperability', 'Data Portability', 'Data Sovereignty', 'Data Residency', 'Data Localization', 'Data Transfer', 'Data Exchange', 'Data Sharing', 'Data Collaboration', 'Data Partnership', 'Data Alliance', 'Data Consortium', 'Data Trust', 'Data Cooperative', 'Data Union', 'Data Federation', 'Data Federation Frameworks', 'Data Federation Standards', 'Data Federation Protocols', 'Data Federation APIs', 'Data Federation SDKs', 'Data Federation Libraries', 'Data Federation Frameworks', 'Data Federation Standards', 'Data Federation Protocols', 'Data Federation APIs', 'Data Federation SDKs', 'Data Federation Libraries'.



Постановка завдання до дипломної роботи

- **Задача 1.** Відновлення лінійної регресії методами байєсівського аналізу (в залежності від апріорних знань) та порівняння з класичним алгоритмом МНК.
- **Задача 2.** Класифікація текстів за допомогою Naïve Bayes classifier та Maximum entropy classifier та порівняння із класичним алгоритмом Support vector machine (SVM).
- **Задача 3.** Автоматичне знаходження розмірності підпростору в задачі зменшення розмірності даних за допомогою Байєсівського PCA та порівняння з класичним алгоритмом PCA.

Задача 1. Відновлення лінійної регресії

- **Постановка задачі:** Нехай є дані, отримані в результаті гіпотетичного експерименту, які лінійно залежать від деякої величини X . Дані надходять з шумом, дисперсія якого невідома. Необхідно знайти коефіцієнти лінійної залежності.

- **Вхідні дані:**

$$x_1 \dots x_m \in [0; 10], \quad m = 20$$

$$y = a \cdot x + b + \varepsilon, \quad a = 8.7, b = 4.4$$

ε – гаусівський стандартний білий шум

Задача 1. Байєсівський аналіз

- Функція правдоподібності:

$$f(X|a, b, \sigma) = \prod_{i=1}^m \mathcal{N}(a \cdot x_i + b, \sigma^2), \sigma > 0$$

- Апріорні розподіли параметрів:

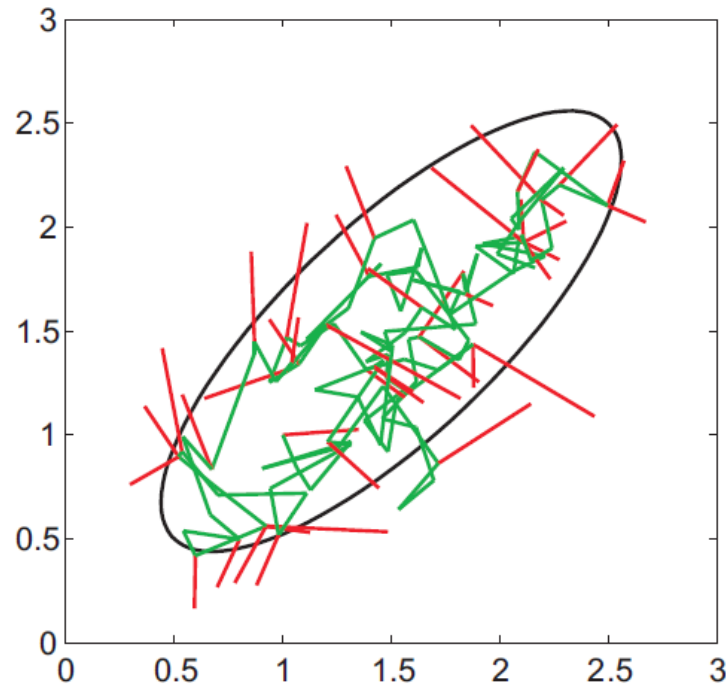
$$a \sim U(\Delta_a), b \sim U(\Delta_b), \sigma \sim U(\Delta_\sigma) \text{ для деяких } \Delta_a, \Delta_b \subset \mathbb{R} \text{ та } \Delta_\sigma \subset \mathbb{R}_+$$

- Оцінки $\hat{a} = \mathbb{E}_p[a]$, $\hat{b} = \mathbb{E}_p[b]$, де:

$$p(a, b, \sigma|X) \propto f(X|a, b, \sigma) \cdot \pi(a, b, \sigma) = f(X|a, b, \sigma) \cdot \pi(a) \cdot \pi(b) \cdot \pi(\sigma)$$

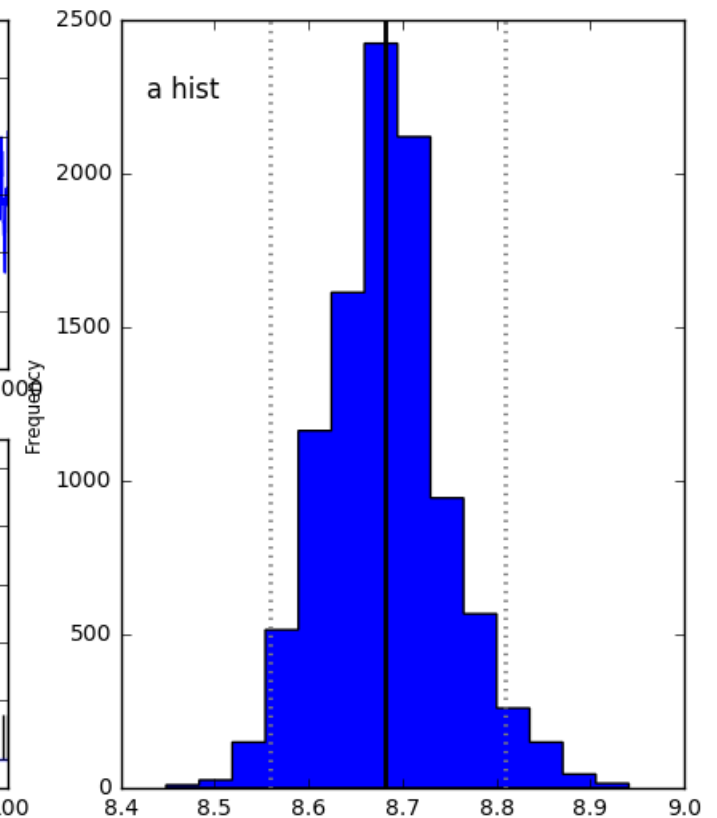
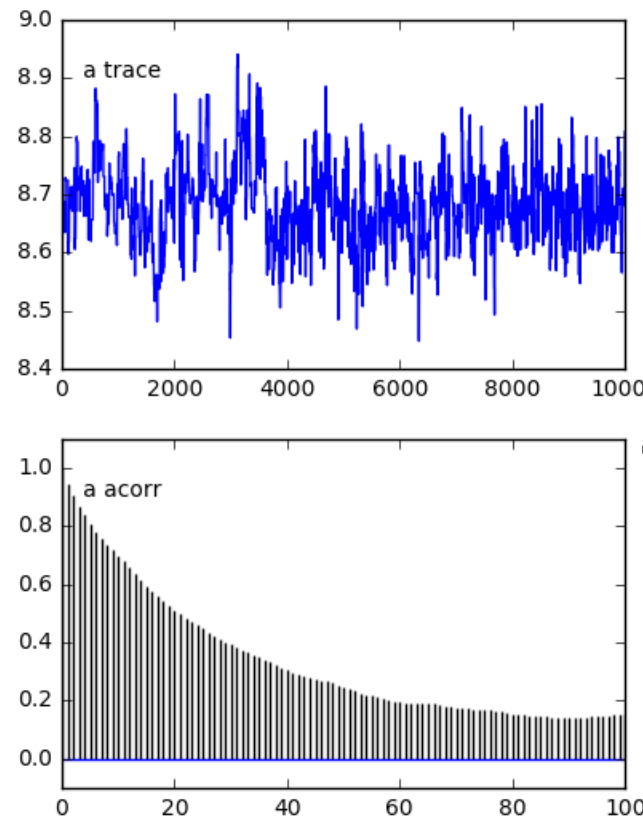
Задача 1. Моделювання апостеріорного розподілу (MCMC)

МСМС за схемою
Метрополіса-Хастінгса:



Параметри:

- кількість ітерацій: 15000;
- період розігріву: 5000.



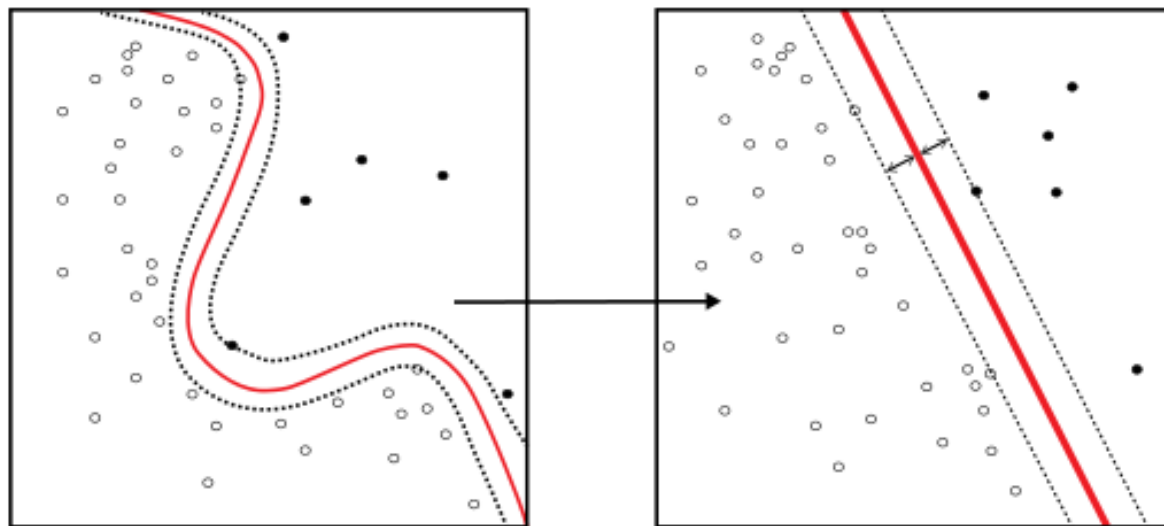
Задача 1. Результати

$\hat{a}_{LSM}$	$\hat{b}_{LSM}$	RMSE
8.6781	4.7776	0.2761

a, b	σ	$\hat{a}_{PM}$	$\hat{b}_{PM}$	RMSE
$\sim U(-10^4, 10^4)$	$\sim U(0, 10^5)$	8.6081	5.1004	0.3688
	$\sim U(0, 10^3)$	8.7295	4.4979	0.2609
	$\sim U(0, 10)$	8.6619	4.8514	0.2855
	$\sim U(0, 2)$	8.6650	4.9426	0.3827
$\sim U(0, 100)$	$\sim U(0, 10^5)$	8.7224	4.6078	0.3267
	$\sim U(0, 10^3)$	8.6851	4.6268	0.1589
	$\sim U(0, 10)$	8.6551	4.7836	0.2096
	$\sim U(0, 2)$	8.7026	4.5636	0.1768
$\sim U(0, 10)$	$\sim U(0, 10^5)$	8.7189	4.4796	0.1830
	$\sim U(0, 10^3)$	8.6818	4.7015	0.2176
	$\sim U(0, 10)$	8.6992	4.5957	0.1917
	$\sim U(0, 2)$	8.6765	4.6193	0.1244

Задача 2. Класифікація текстів

- Класифікація текстів:
 - спам-фільтри
 - класифікація пошти на групи
 - автоматична ідентифікація мови, жанру, настрою автора та ін.
- До сих пір не існує універсального алгоритму. Один з лідерів – SVM. Головні недоліки:
 - лінійність;
 - складно інтерпретувати деякі параметри;
 - модель не ймовірнісна, складно оцінити $\mathbb{P}(y_i|x)$
- Будуть розглянуті Naïve Bayes classifier та Maximum entropy classifier



Задача 2. Класифікація текстів

- **Постановка задачі:** Нехай D є множина документів та множина категорій C . Нехай також є тренувальний промаркований набір пар «документ-категорія» (d, c) , $d \in D, c \in C$. Необхідно запропонувати алгоритм, який на основі цих даних міг би віднести нові документи до однієї чи декількох категорій з C .
- **Вхідні дані:** відкритий датасет Reuters – 9517 уривків новин за 90 можливими категоріями англійською мовою.
- Побудуємо три бінарних класифікатори для розпізнавання трьох категорій з найбільшою кількістю документів в датасеті – earn, acq та money-fx.

Задача 2. Naïve Bayes classifier

- Моделюється умовний розподіл $\mathbb{P}(c|\mathbf{d})$ – категорії c за умови що спостерігається документ $\mathbf{d}$:

$$\mathbb{P}(c|\mathbf{d}) = \frac{\mathbb{P}(c) \cdot \mathbb{P}(\mathbf{d}|c)}{\mathbb{P}(\mathbf{d})} \propto \mathbb{P}(c) \cdot \mathbb{P}(\mathbf{d}|c)$$

- Припущення умовної незалежності ознак:

$$\forall i \neq j \quad \forall c \in C \quad \mathbb{P}(d_i, d_j|c) = \mathbb{P}(d_i|c) \cdot \mathbb{P}(d_j|c) \quad \Rightarrow \quad \mathbb{P}(\mathbf{d}|c) = \prod_{i=1}^n \mathbb{P}(d_i|c)$$

- Змоделювавши розподіл, клас обирається за оцінкою MAP:

$$\hat{c} = \operatorname{argmax}_{c \in C} \left\{ \mathbb{P}(c) \cdot \prod_{i=1}^n \mathbb{P}(d_i|c) \right\}$$

Задача 2. Принцип максимальної ентропії

- Нехай відома деяка інформація (деякі обмеження) щодо апріорного розподілу.
- Тоді слід обирати апріорний розподіл, який за даних обмежень максимізує ентропію:

$$\mathbb{P}\{X = x_k\} = p_k, \quad k = 1, 2, \dots \quad f(x)$$

$$H(X) = -\sum_k p_k \cdot \log(p_k) \quad H(X) = -\int_{-\infty}^{+\infty} f(x) \log(f(x)) dx$$

Задача 2. Maximum entropy classifier

- Моделюється $\mathbb{P}(c|d)$.
- $f_i(d, c) \in \mathbb{R}$ – набір ознак документів. Обмеження:

$$\forall i \quad \frac{1}{|D|} \sum_{d \in D} f_i(d, c(d)) = \mathbb{E}_{c|d}[f_i(d, c)]$$

- Можна показати, що за цих обмежень розподіл матиме наступну форму:

$$\mathbb{P}_{\Lambda}(c|d) \propto \exp \left(\sum_i \lambda_i f_i(d, c) \right),$$

λ_i – параметри, які необхідно оцінити.

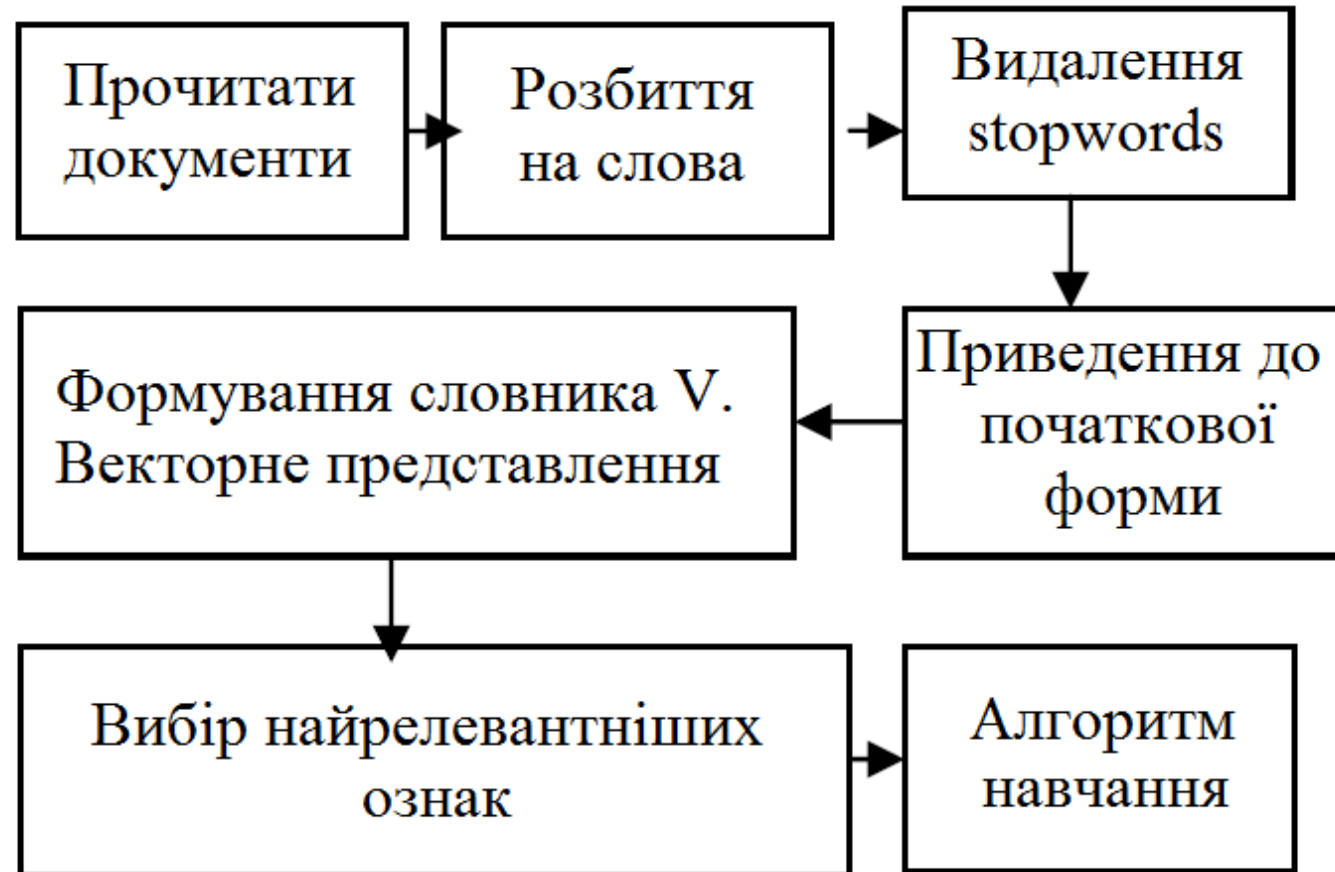
Задача 2. Maximum entropy classifier

- Можна показати, що задача максимізації ентропії в даному випадку еквівалентна максимізації (логарифмічної) правдоподібності:

$$\ell(\Lambda|D) = \log \left(\prod_{d \in D} \mathbb{P}_{\Lambda}(c(d)|d) \right) = \sum_{d \in D} \log \sum_{c \in C} \exp \sum_i \lambda_i f_i(d, c(d)) \rightarrow \max_{\Lambda}$$

- Для максимізації $\ell(\Lambda|D)$ використано ітераційний алгоритм Improved Iterative Scaling (IIS), ідея якого базується на опуклості вгору даної функції. Умова зупинки алгоритму $|\ell_{n+1} - \ell_n| < 10^{-6}$.

Задача 2. Алгоритм виділення ознак з документів



Задача 2. Постановка експерименту

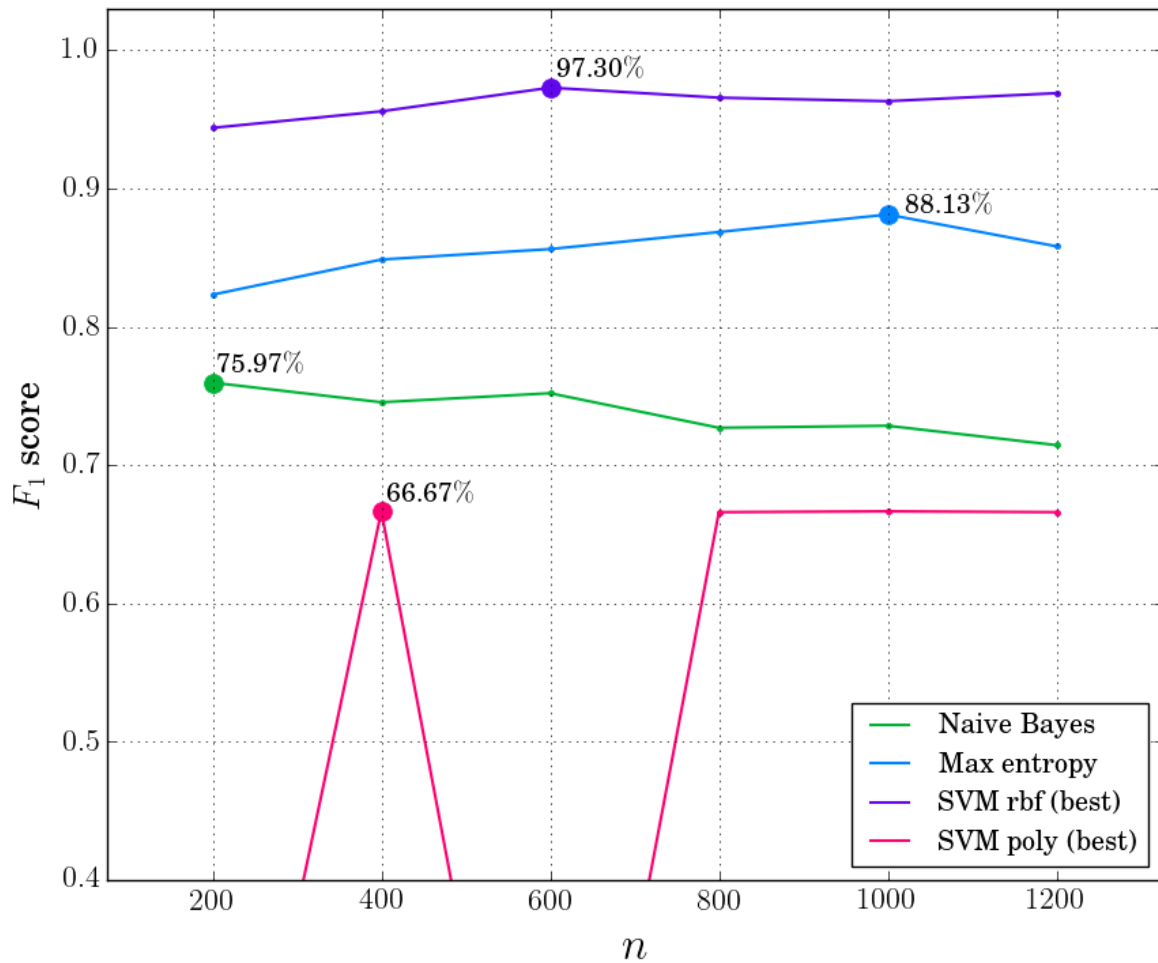
- Кількість ознак $n \in \{200, 400, 600, 800, 1000, 1200\}$.
- SVM запускався з наступними ядрами:
 - гаусівське (rbf): $k(\mathbf{x}_i, \mathbf{x}_j) = \exp(-\gamma \cdot \|\mathbf{x}_i - \mathbf{x}_j\|^2)$ для $\gamma \in \{0.1, 0.5, 1.0, 2.0, 5.0, 10, 20, 50\}$;
 - поліноміальне: $k(\mathbf{x}_i, \mathbf{x}_j) = (\mathbf{x}_i \cdot \mathbf{x}_j)^d$ для $d \in \{1, 2, 3, 4, 5\}$.
- Вибірка = Тренувальна (80%) : cross-validation (10%) : тестова (10%)
- Метрика – оцінка F_1 :

$$F_1 = \frac{2PR}{P + R}, P(\text{precision}) = \frac{\text{кількість документів, що правильно передбачені як 1}}{\text{кількість документів, що передбачені як 1}},$$

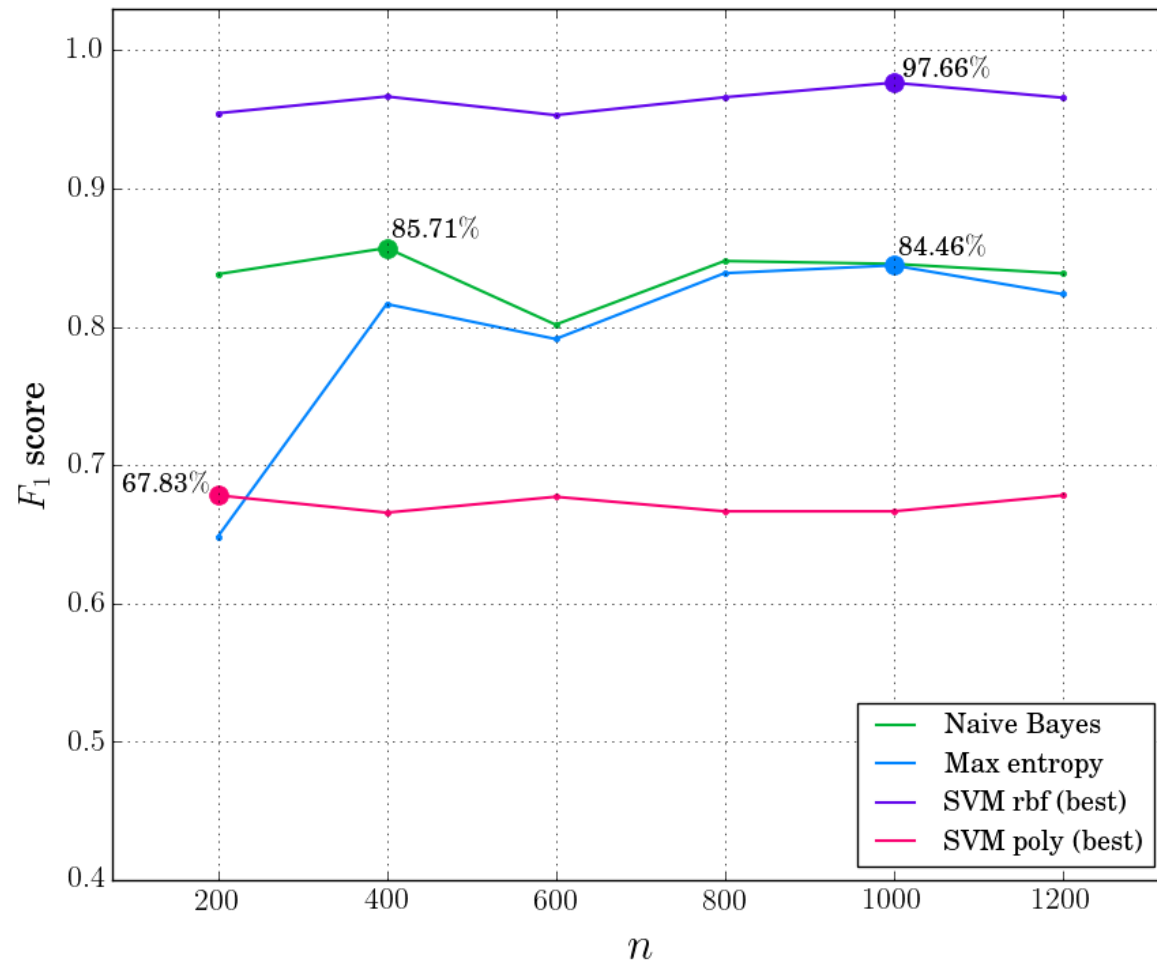
$$R(\text{recall}) = \frac{\text{кількість документів, що правильно передбачені як 1}}{\text{кількість документів, що дійсно належать до 1}}$$

Задача 2. Результати

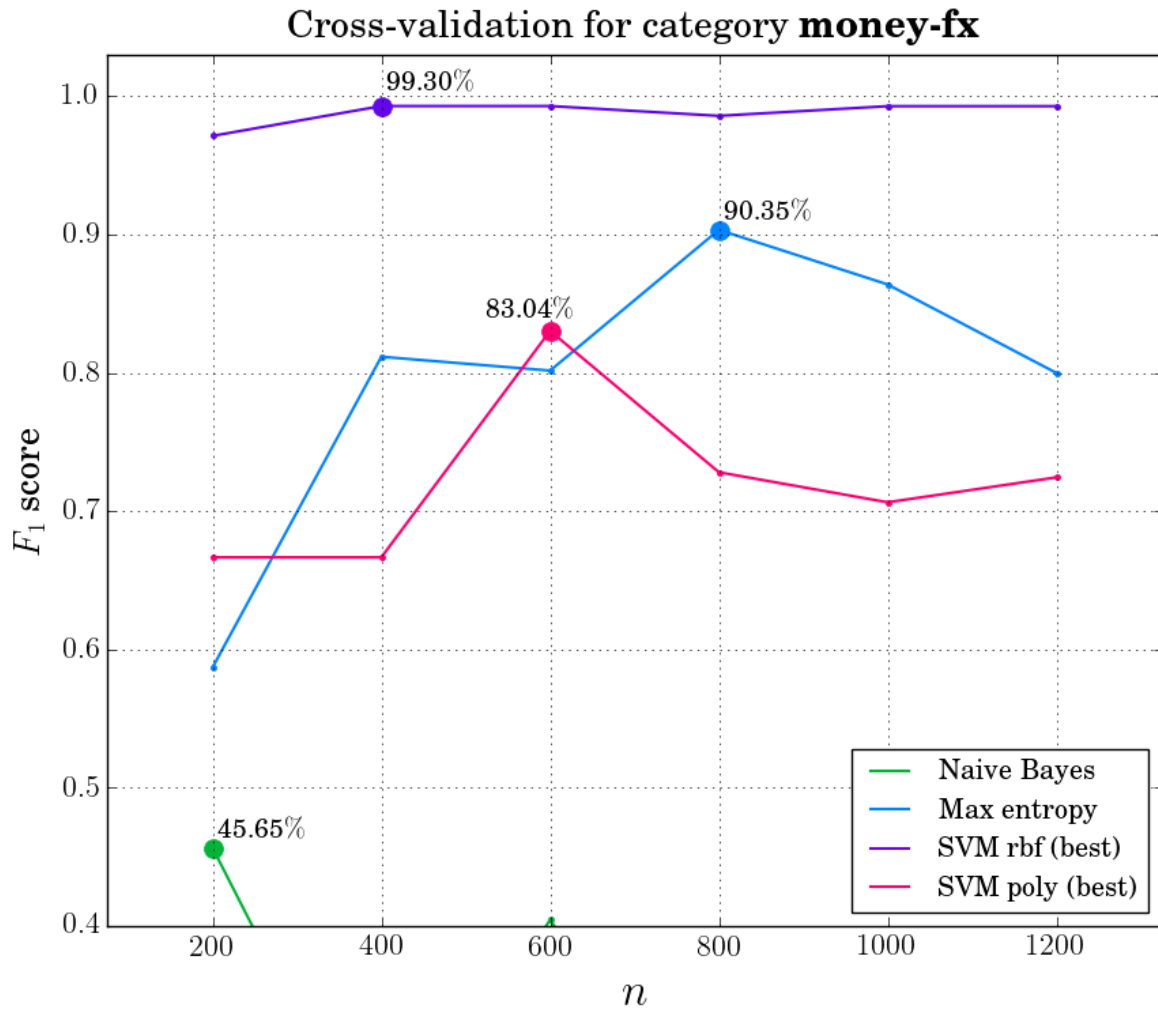
Cross-validation for category **earn**



Cross-validation for category **acq**



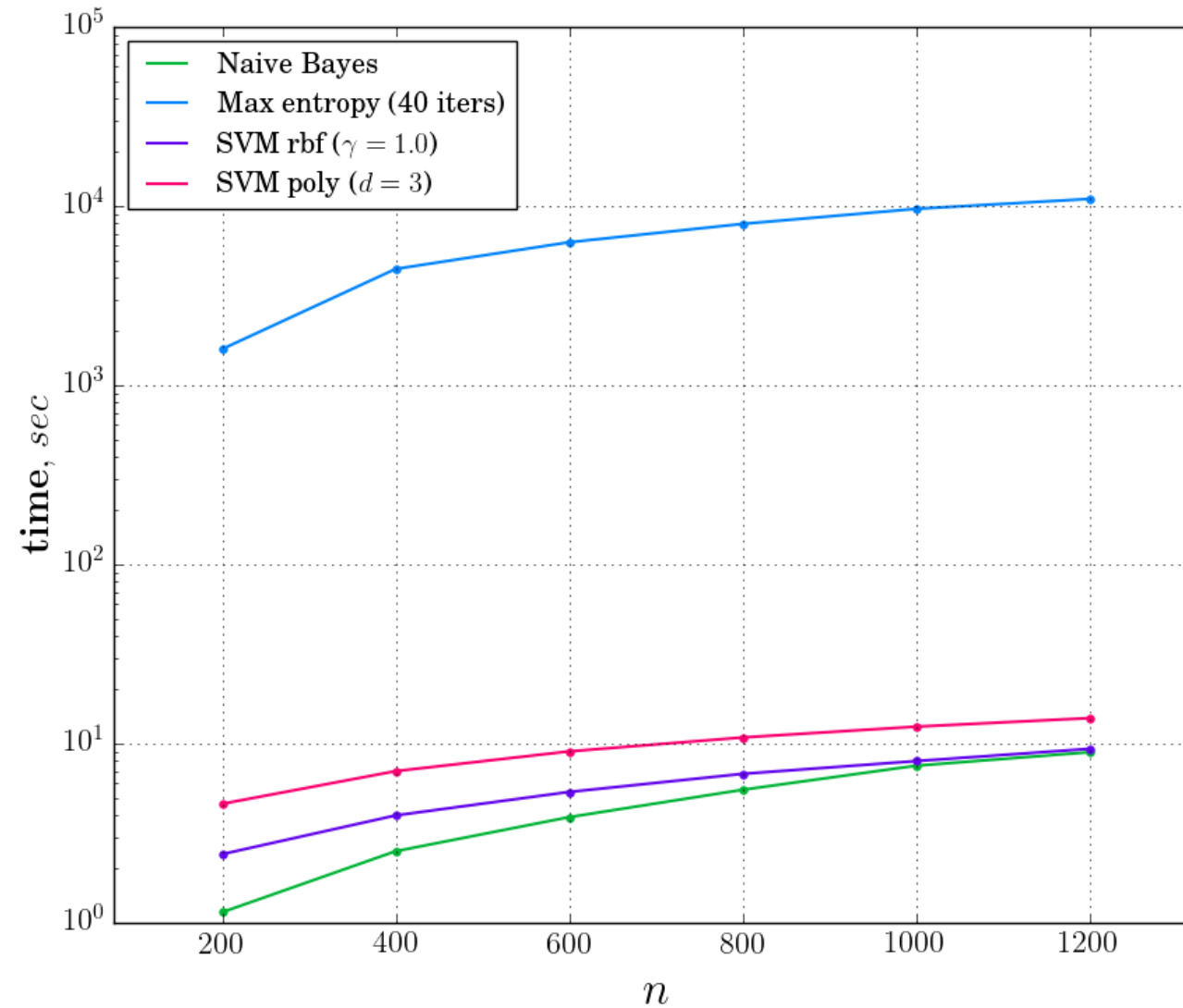
Задача 2. Результати



Оцінки F_1 для тестової вибірки для кращих моделей з cross-validation

Категорія	NBC	MEC	SVM
earn	77.56%	88.86%	95.95%
acq	81.55%	82.92%	95.10%
money-fx	42.22%	86.27%	97.90%

Задача 2. Час роботи алгоритмів



Задача 3. Зменшення розмірності даних

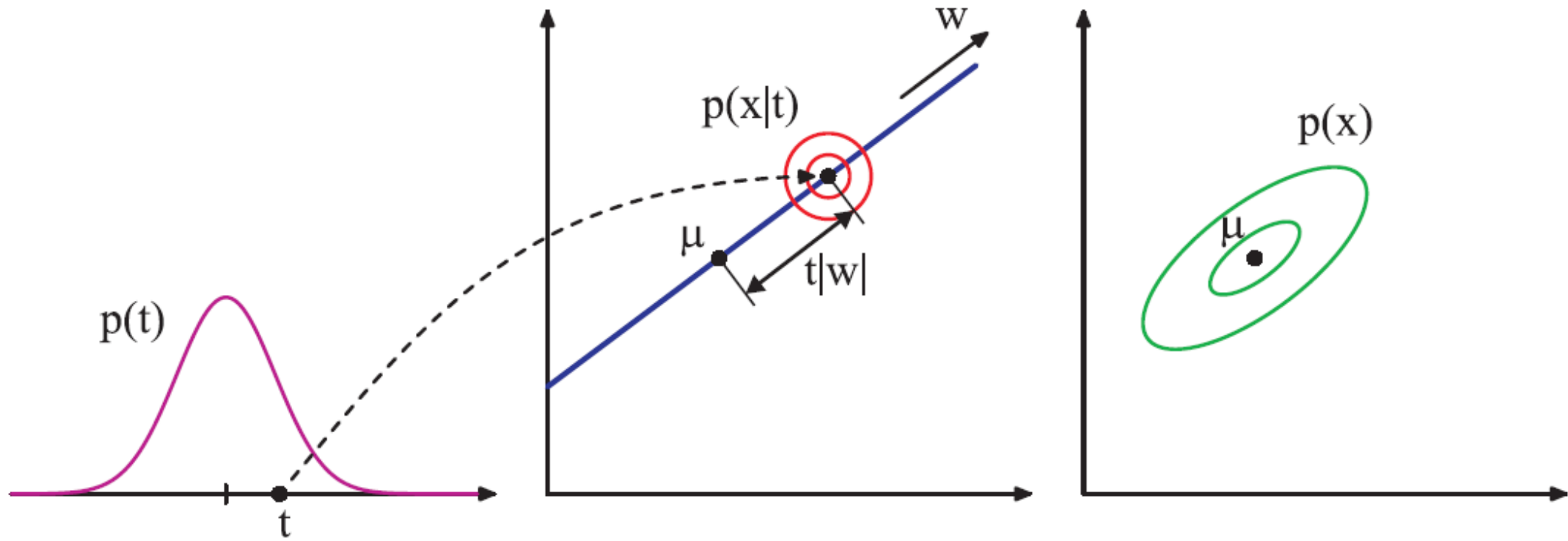
- Задача зменшення розмірності даних:
 - зменшення обчислювальних витрат
 - стиснення даних для ефективного їх зберігання
 - візуалізація даних
 - виявлення нових (прихованих) ознак в даних та ін.
- Класичний алгоритм – PCA, що полягає у пошуку ортонабору векторів $w_1 \dots w_d$ в просторі $\mathbb{R}^n$ ($d < n$) – *головних компонент*, що дисперсія спроектованих даних на їх лінійну оболонку буде максимальною $\Leftrightarrow$ середньоквадратична похибка реконструкції буде мінімальною. Головні недоліки:
 - до уваги беруться тільки середнє та коваріація
 - міра інформативності – дисперсія
 - лінійність
 - трудомісткі операції
 - необхідність задання параметра d
- Будуть розглянуті ймовірнісні моделі PCA.

Задача 3. Зменшення розмірності даних

- **Постановка задачі:** Нехай є набір спостережуваних даних $X = \{x_1 \dots x_m\} \in \mathbb{R}^n$. Дослідимо задачу автоматичного обрання *ефективної* розмірності підпростору для їх представлення, а також дослідимо питання зменшення похибки реконструкції в цьому підпросторі порівняно із класичним алгоритмом PCA.
- **Ймовірнісна модель PCA (PPCA):**
 - набір прихованих змінних $T = \{t_1 \dots t_m\} \in \mathbb{R}^d$, $d < n$, d – бажана розмірність;
 - $p(x|t) = \mathcal{N}(x|Wt + \mu, \sigma^2 I)$, $p(t) = \mathcal{N}(t|0, I) \implies \implies p(x) = \mathcal{N}(x|\mu, C = WW^T + \sigma^2 I \in \mathbb{R}^{n \times n})$
 - параметри моделі W, μ, σ можна знайти методом максимальної правдоподібності (аналітично, або чисельно за допомогою EM-алгоритму);
 - знайшовши параметри, представлення для x є мат сподіванням $p(t|x)$:
$$\mathbb{E}_{t|x} t = M^{-1} W^T (x - \mu), \quad M = W^T W + \sigma^2 I \in \mathbb{R}^{d \times d}$$

Задача 3. Модель РРСА

Генерація об'єктів у моделі РРСА для $n = 2$ для $d = 1$:



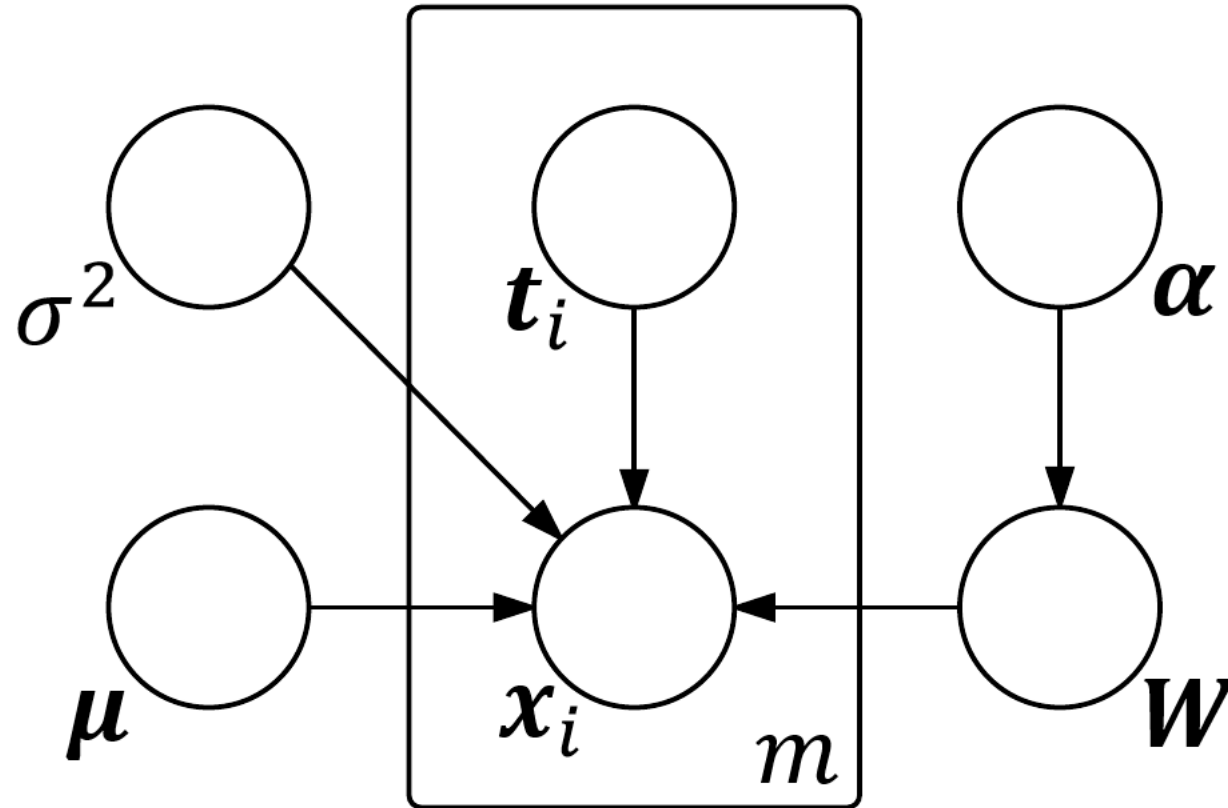
Задача 3. Модель ВРСА

Байєсівська модель РСА (ВРСА):

- $d := n - 1$
- $p(\mathbf{W}|\boldsymbol{\alpha}) = \prod_{i=1}^{n-1} \left(\frac{\alpha_i}{2\pi}\right)^{n/2} \exp\left\{-\frac{1}{2}\alpha_i\|\mathbf{w}_i\|^2\right\}$, гіперпараметри α_i контролюють обернені дисперсії відповідних $\mathbf{w}_i$: $\alpha_i \rightarrow +\infty \Rightarrow \mathbf{w}_i \rightarrow \mathbf{0}$ – можна відкинути
- Параметри знаходимо методом максимальної правдоподібності за допомогою ЕМ-алгоритму:

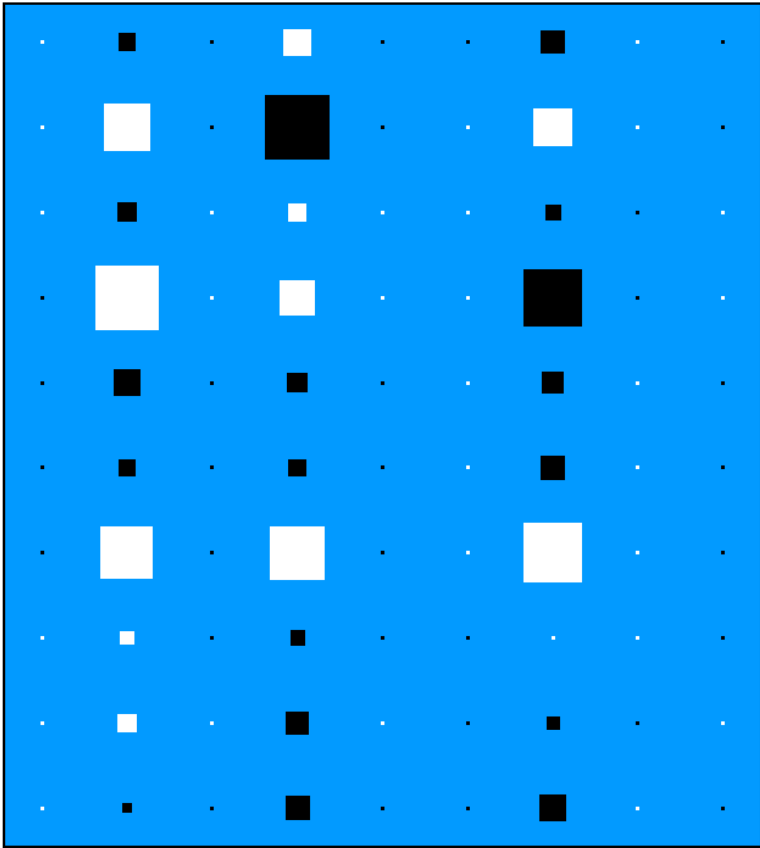
$$p(X|\boldsymbol{\mu}, \sigma, \boldsymbol{\alpha}) = \int p(X|\mathbf{W}, \boldsymbol{\mu}, \sigma) p(\mathbf{W}|\boldsymbol{\alpha}) d\mathbf{W} \rightarrow \max_{\boldsymbol{\mu}, \sigma, \boldsymbol{\alpha}}$$

Задача 3. Графічне представлення моделі ВРСА



Задача 3. Опис даних та результати

Вхідні дані: $x_1 \dots x_{300} \sim N(\mathbf{0}, \text{diag}\{\sigma_1^2 \dots \sigma_{10}^2\})$, $\sigma_i = 1, i \in \{1, 2, 3\}$ та $\sigma_i = 0.5$ для решти i



Матриця W , обчислена за методом ВРСА

Алгоритм	MRE
PCA, $d = 5$	1.2812
PCA, $d = 6$	0.9654
PCA, $d = 7$	0.7196
PCA, $d = 8$	0.5593
PCA, $d = 9$	0.2423
ВРСА ($d_{eff} = 3$)	0.1851

Порівняння ВРСА із PCA за *середньою похибкою реконструкції* (MRE) – $\frac{1}{m} \sum_{i=1}^m \|x_i - \hat{x}_i\|^2$

Висновки та подальші дослідження

Висновки:

- На прикладі лінійної регресії досліджена залежність якості оцінок за допомогою байєсівського аналізу від ступеня точності апріорних знань
- Для задачі класифікації текстів порівняні алгоритми NBC, MEC із класичним SVM
- Автоматична ідентифікація розмірності підпростору в задачі зменшення даних, зменшення похибки реконструкції порівняно із PCA

Подальші дослідження:

- Для розглянутих задач більш ефективні алгоритми чи їх модифікації
- Розширити коло задач (нейронні мережі)

Дякую за увагу!