

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ
“КИЇВСЬКИЙ ПОЛІТЕХНІЧНИЙ ІНСТИТУТ ІМ. ІГОРЯ СІКОРСЬКОГО”
ІНСТИТУТ ПРИКЛАДНОГО СИСТЕМНОГО АНАЛІЗУ

НЕЙРОННІ МЕРЕЖІ ДЛЯ РОЗПІЗНАВАННЯ ЕМОЦІЙ У ВІДЕО ТА АУДІО ПОТОКАХ

Роботу виконали:

Антон Щербина

Вадим Коршунов

Науковий керівник:

Юрій Тимошенко

- ▶ Неперервна діагностика емоційного стану
- ▶ Додаткова інформація для маркетингових досліджень та рекомендаційних систем
- ▶ Адаптація інтерфейсів взаємодії з комп'ютером

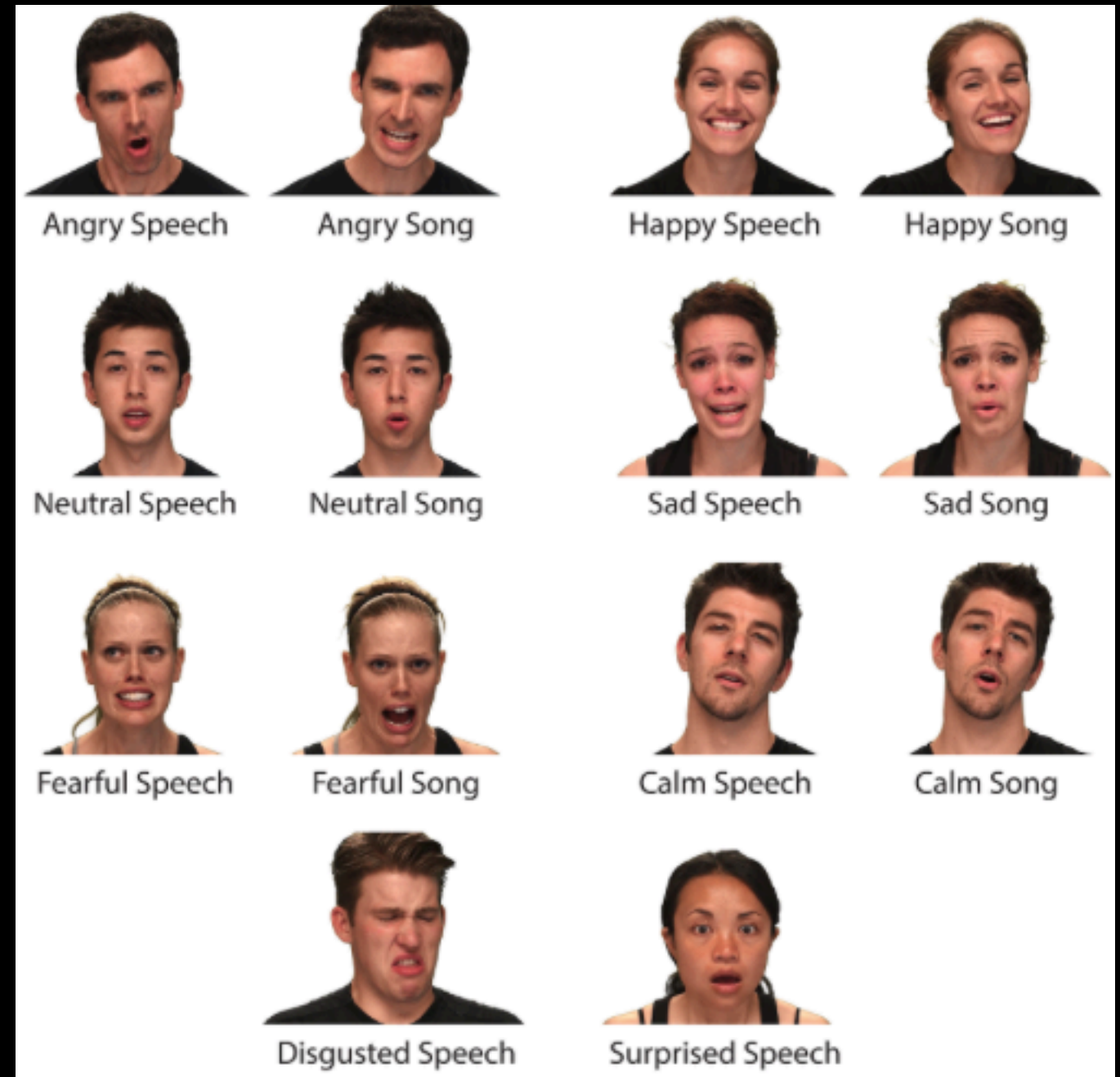
- ▶ Мета. Дослідити методи глибокого навчання для вирішення задачі класифікації емоції з відео та аудіо в режимі реального часу, імплементувати та протестувати найкращі моделі.
- ▶ Об'єкт. Розпізнавання людських емоцій за відео та аудіо в режимі реального часу.
- ▶ Предмет. Методи глибокого навчання, які дозволяють вирішити проблему розпізнавання емоцій.

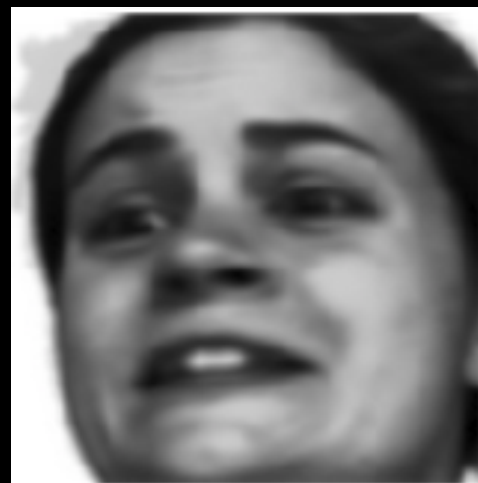
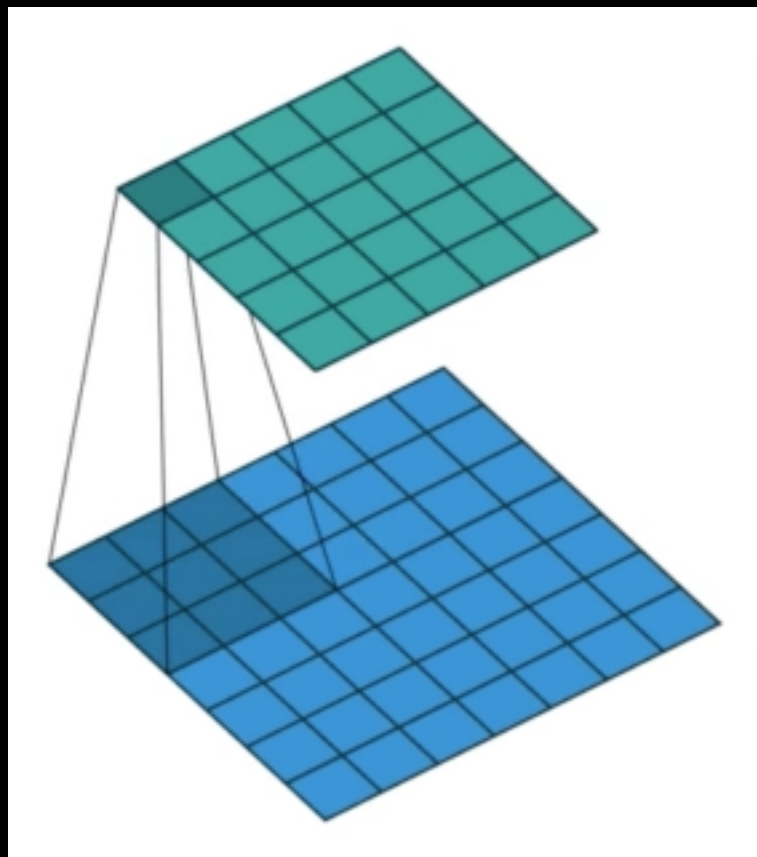
- ▶ Ручне виділення ознак
- ▶ Необхідне знання доменної області
- ▶ Низька узагальнююча здатність

- ▶ Не потребують ручного виділення ознак
- ▶ Дозволяють працювати з даними великої розмірності
- ▶ Можуть навчатись end-to-end

- ▶ Збір даних та їх передобробка
- ▶ Побудова архітектур моделей
- ▶ Оптимізація параметрів моделей
- ▶ Оцінка генералізуючої здатності кожної з моделей
- ▶ Ансамблювання моделей

- ▶ 4900 високоякісних відео та аудіо файлів
- ▶ 24 актори
- ▶ 8 емоцій
- ▶ 2 типи інтенсивності
- ▶ 2 типи мовлення
- ▶ 2 типи повідомлення



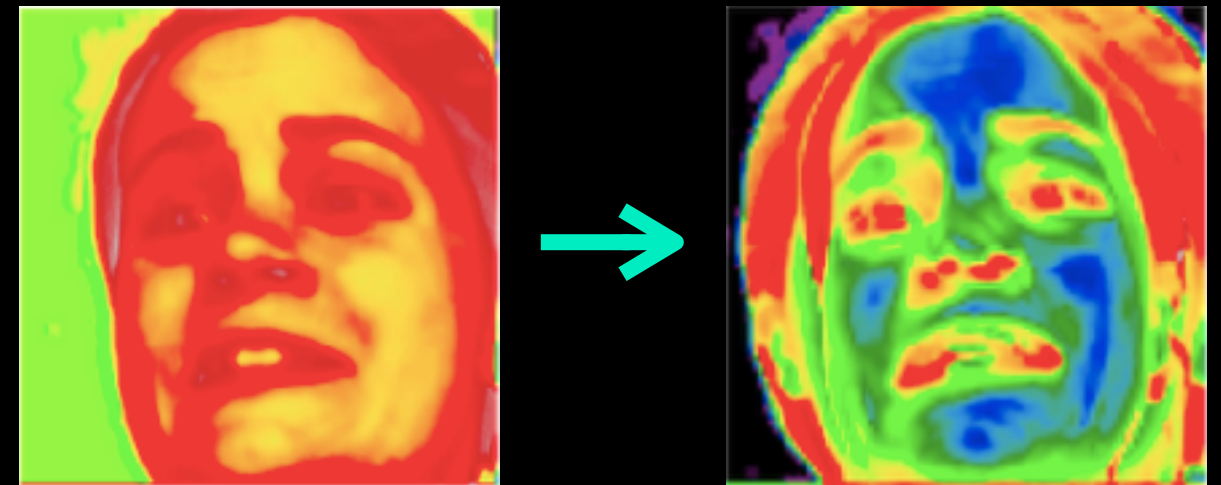
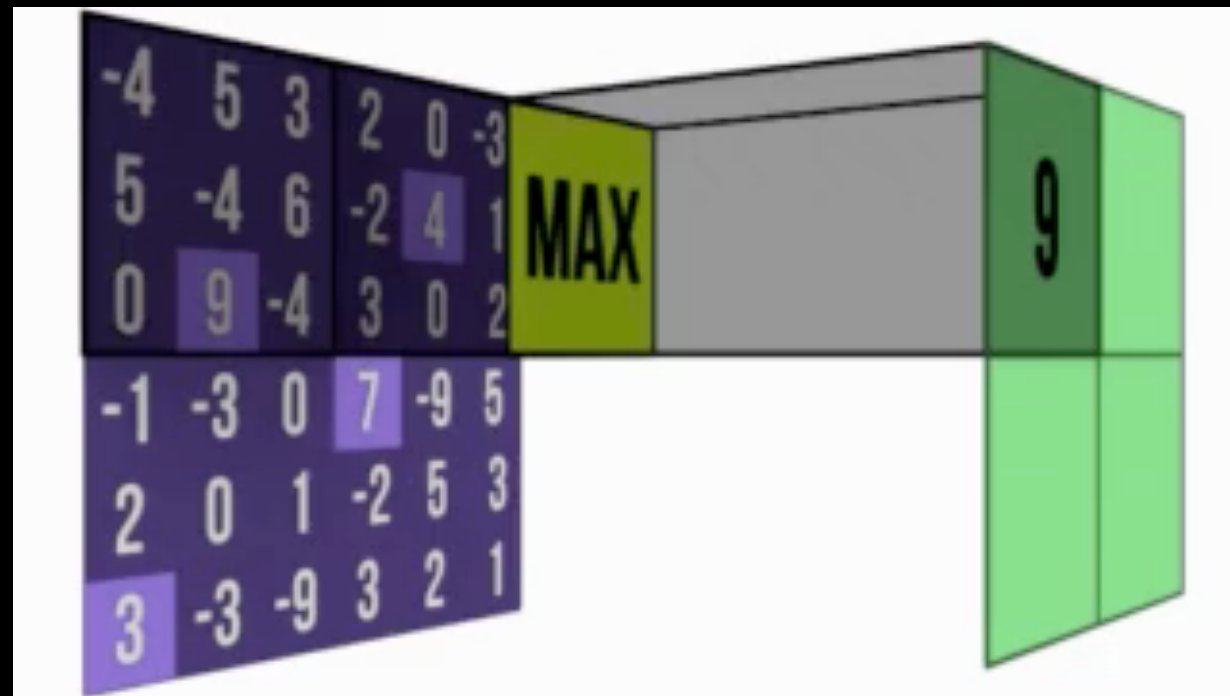


$$Output[i][j] = \sum_{k=1}^m \sum_{l=1}^n Image(i+k-1, j+l-1) * Kernel(k, l)$$

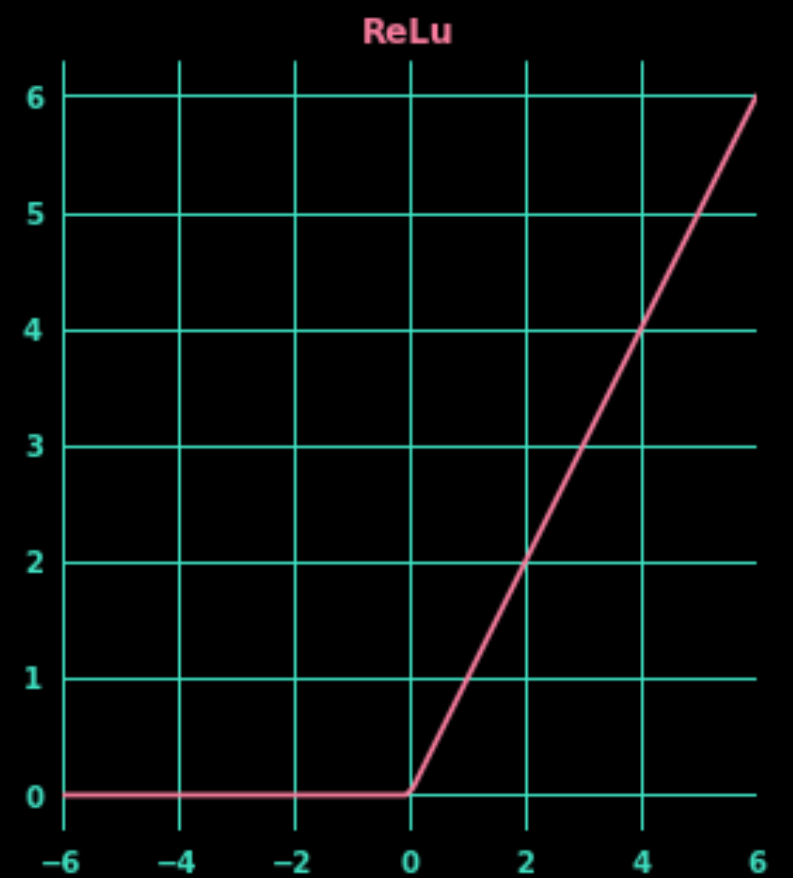
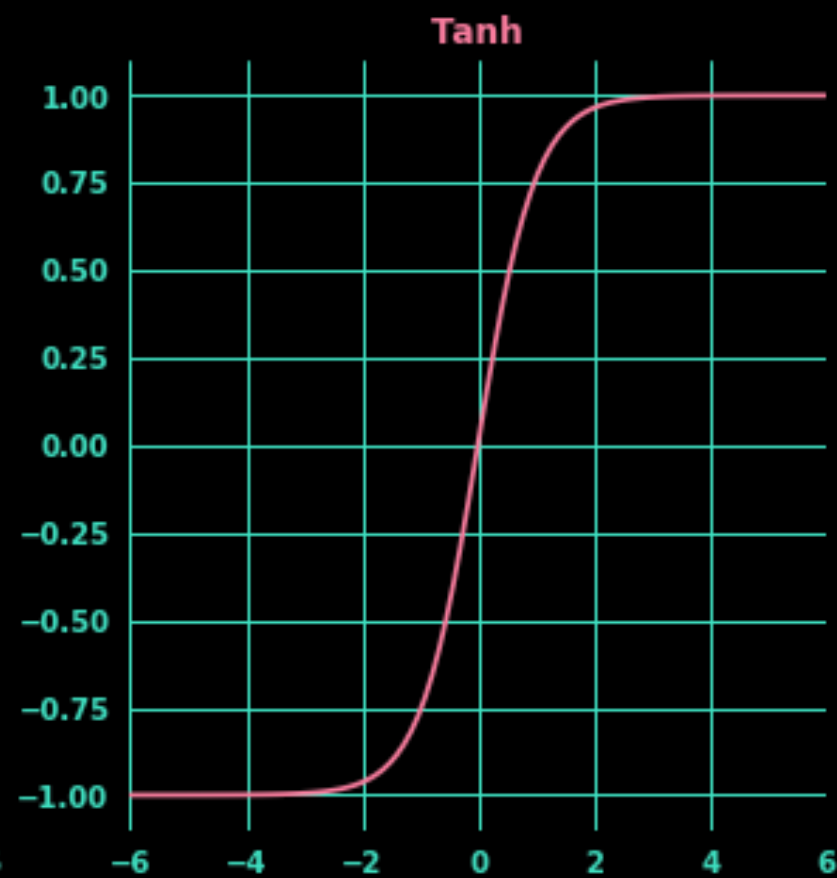
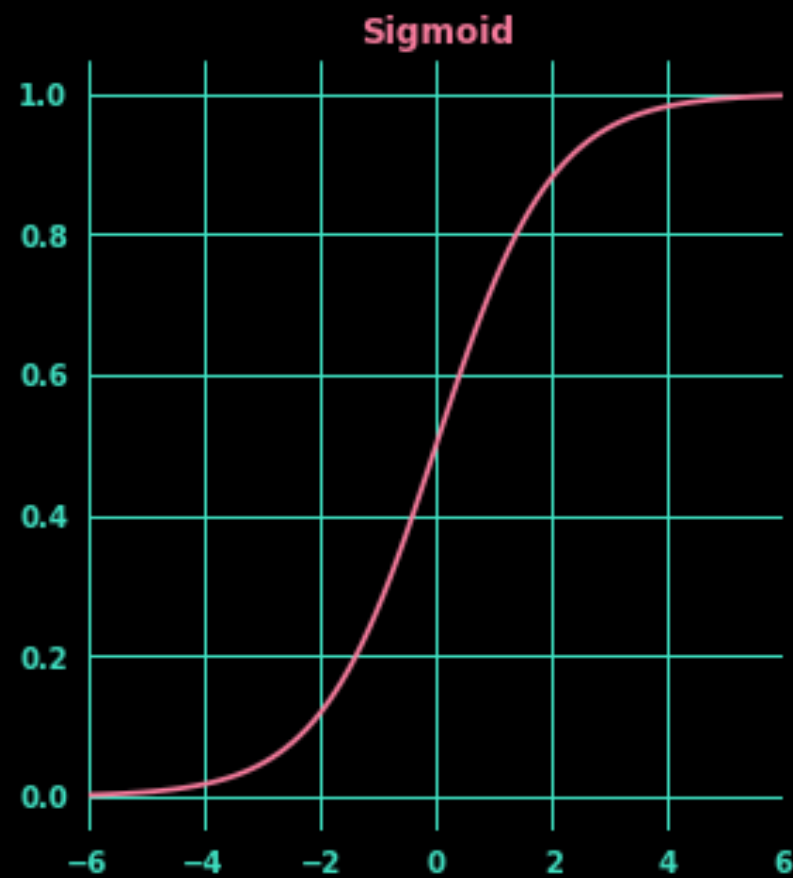
$$i = 1, \dots, M - m + 1, j = 1 \dots N - n + 1$$

MAX POOLING

9



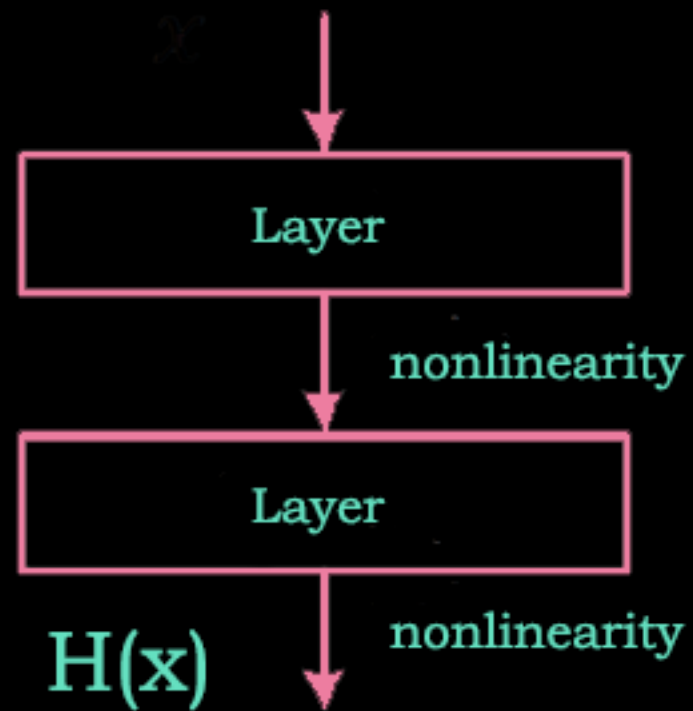
$$Output[i][j] = \max\{Image(i+k-1, j+l-1), k = 1 \dots m, l = 1 \dots n, i = 1, \dots, M-m+1, j = 1 \dots N-n+1\}$$



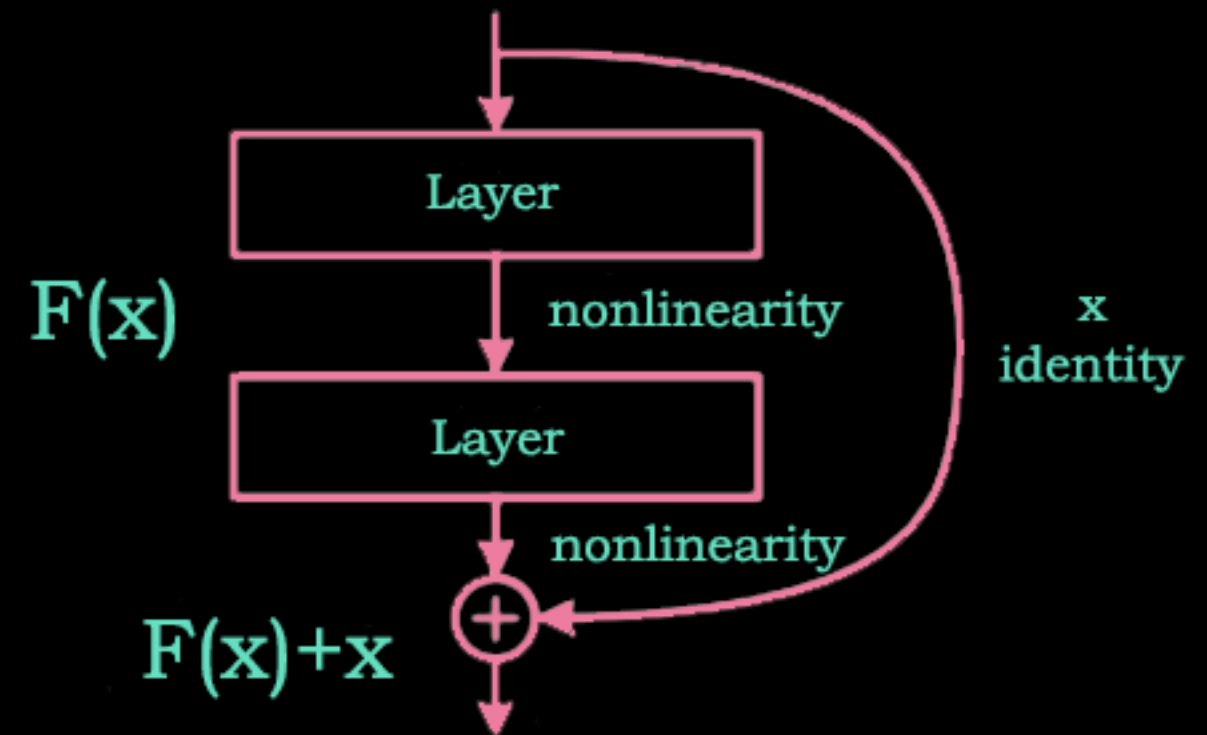
$$f(x) = \frac{1}{1 + e^{-x}}$$

$$f(x) = \frac{2}{1 + e^{-2x}} - 1$$

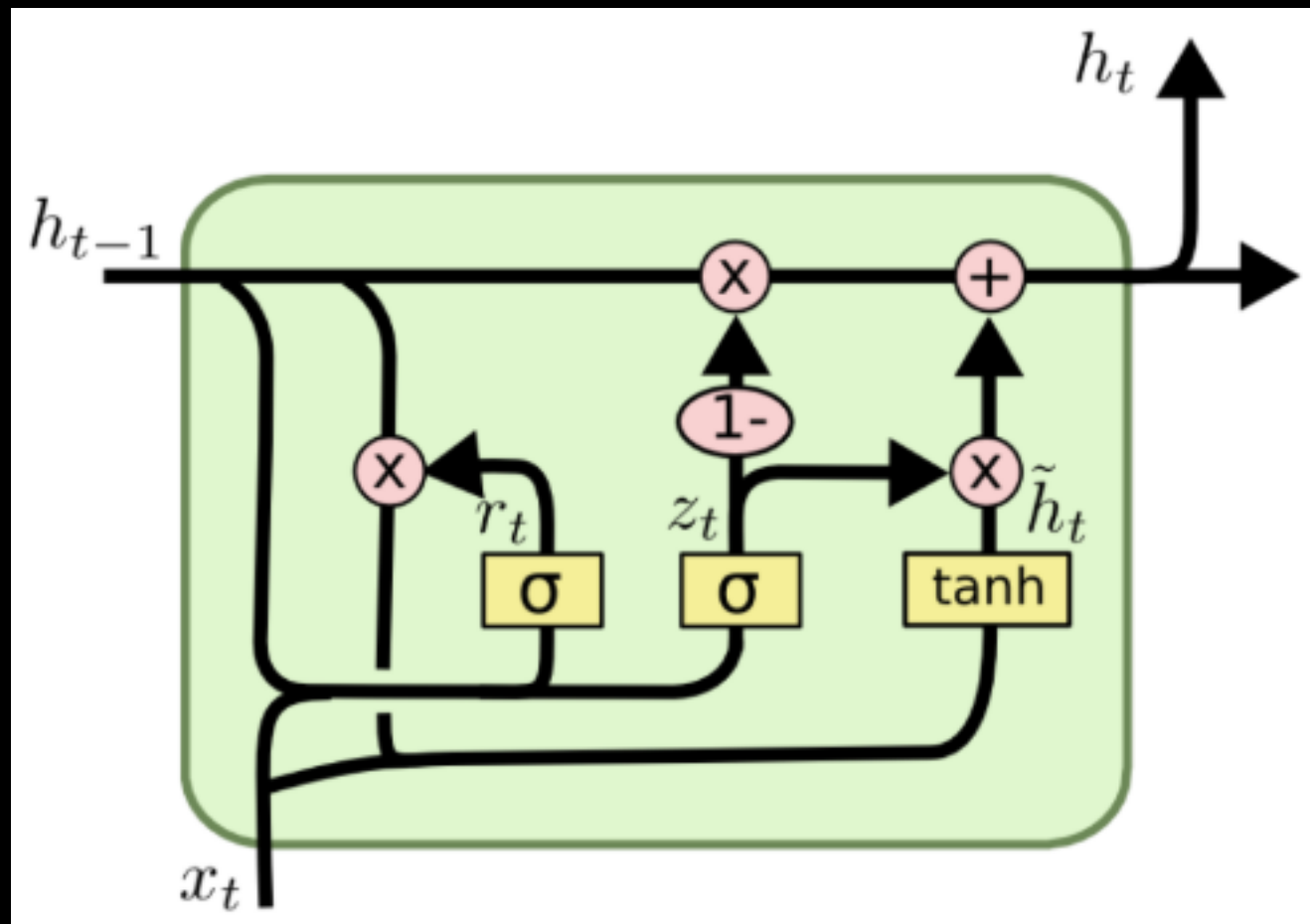
$$f(x) = \max(0, x)$$



Звичайний блок



Залишковий блок



$$z_t = \sigma(W_z[h_{t-1}, x_t])$$

$$r_t = \sigma(W_r[h_{t-1}, x_t])$$

$$\hat{h}_t = \tanh(W[h_{t-1}, x_t] * r_t)$$

$$h_t = (1 - z_t) * h_{t-1} + z_t * \hat{h}_t$$

<http://colah.github.io/posts/2015-08-Understanding-LSTMs/>

$$\textit{softmax}(z) = \begin{pmatrix} \frac{e^{z_1}}{\sum_{i=1}^n e^{z_i}} \\ \vdots \\ \frac{e^{z_n}}{\sum_{i=1}^n e^{z_i}} \end{pmatrix} \quad \begin{pmatrix} 1.2 \\ 0.9 \\ 0.4 \end{pmatrix} \xrightarrow{\textit{softmax}} \begin{pmatrix} 0.46 \\ 0.34 \\ 0.20 \end{pmatrix}$$

- ▶ Нехай є $\langle x_1, l_1 \rangle, \dots, \langle x_N, l_N \rangle$ - навчальні приклади, де x - вхідний елемент, l - його відповідний клас
- ▶ Відомий модельований розподіл класів для кожного вхідного елементу та його реальний розподіл класів:

$$\forall x_i : Model(x_i; \theta) = (p_1(x_i; \theta), \dots, p_n(x_i; \theta))$$

$$\forall x_i : Real(x_i; \theta) = (0_1, \dots, 1_{l_i}, \dots, 0_n)$$

- ▶ Як наблизити розподіл (q) іншим модельованим (p)?

- ▶ Нехай є модельований розподіл p та справжній розподіл q
- ▶ Введемо міру "відстані" між розподілами - перехресну ентропію:

$$H(p, q) = - \sum_{i=1}^n q_i \log(p_i) = H(q) + D_{KL}(q || p)$$

- ▶ На основі перехресної ентропії вводиться функція втрат для багатокласової класифікації на всьому наборі даних:

$$L(\theta) = \frac{1}{N} \sum_{i=1}^N \sum_{j=1}^n q_j(x_i) \log(p_j(x_i; \theta)) + \alpha || \theta ||^2$$

- ▶ Нехай ξ - дискретна випадкова величина, яка дорівнює рівномірному розподілу навчальних прикладів:

$$\mathbb{P}[\xi = x_i] = \frac{1}{N}$$

- ▶ Береться стохастична апроксимація функції втрат:

$$L^*(\theta) = \frac{1}{M} \sum_{i=1}^M H(p(\xi_i; \theta), q(\xi_i)) + \alpha ||\theta||^2$$

- ▶ Апроксимація функції втрат та її градієнт - незміщені статистичні оцінки:

$$\mathbb{E}L^*(\theta) = L(\theta)$$

$$\mathbb{E} \overrightarrow{\frac{\partial}{\partial \theta}} L^*(\theta) = \overrightarrow{\frac{\partial}{\partial \theta}} L(\theta)$$

$$grad_t \leftarrow grad_{\theta} L_t^*(\theta_{t-1})$$

$$M_t \leftarrow \beta_1 M_{t-1} + (1 - \beta_1) grad_t$$

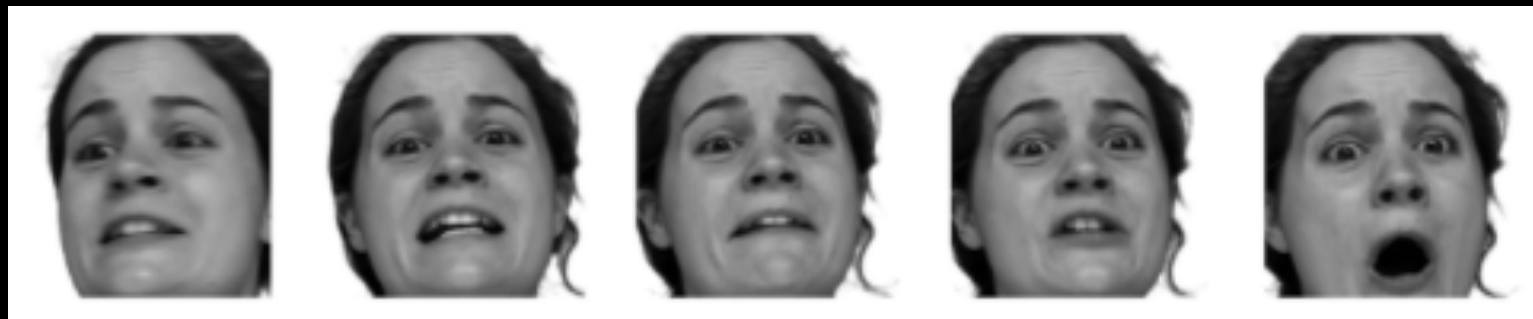
$$V_t \leftarrow \beta_2 V_{t-1} + (1 - \beta_2) grad_t^2$$

$$\tilde{M}_t \leftarrow \frac{M_t}{1 - \beta_1^t}$$

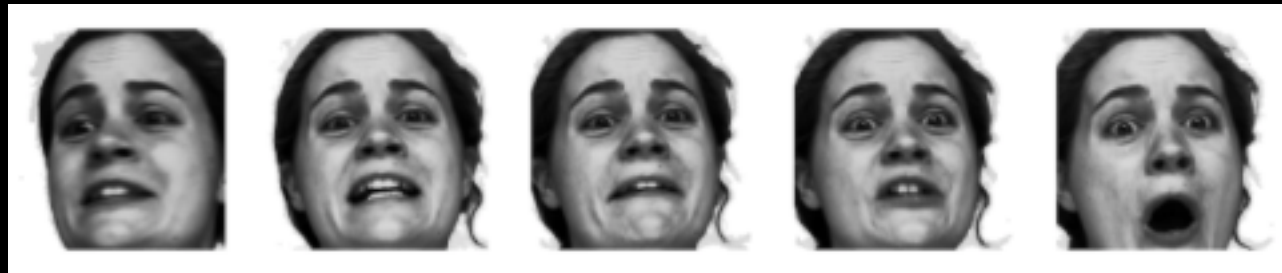
$$\tilde{V}_t \leftarrow \frac{V_t}{1 - \beta_2^t}$$

$$\theta_t \leftarrow \theta_t - \alpha \cdot \frac{\tilde{M}_t}{\tilde{V}_t^{\frac{1}{2}} + \epsilon}$$

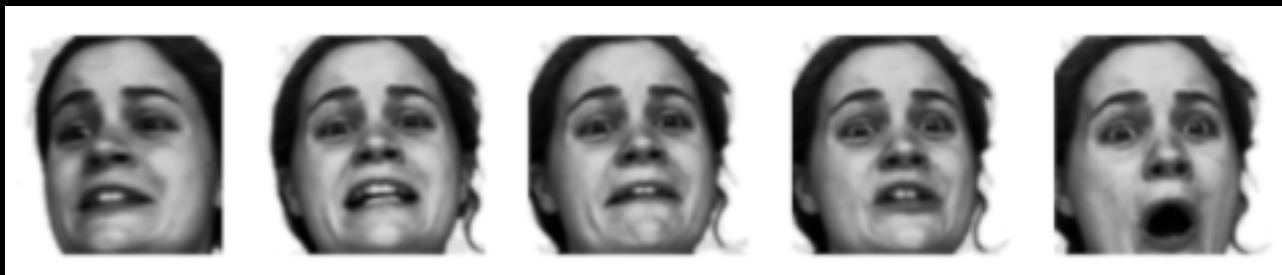
- ▶ Передискретизація вхідного відео
- ▶ Перетворення кольорових кадрів на чорнобілі
- ▶ Виділення облич людей з зображень
- ▶ Масштабування облич до однакового розміру

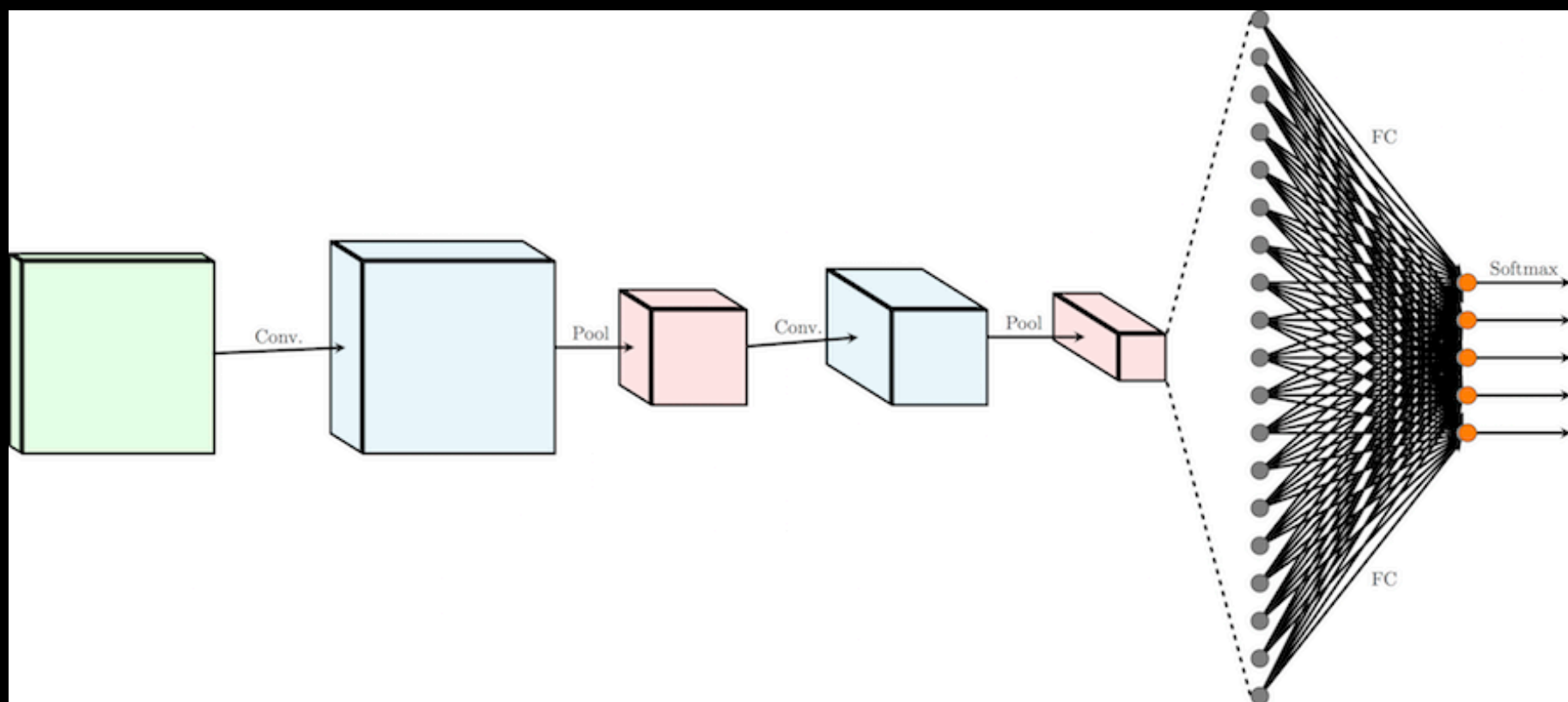


► Вирівнювання гістограми яскравостей

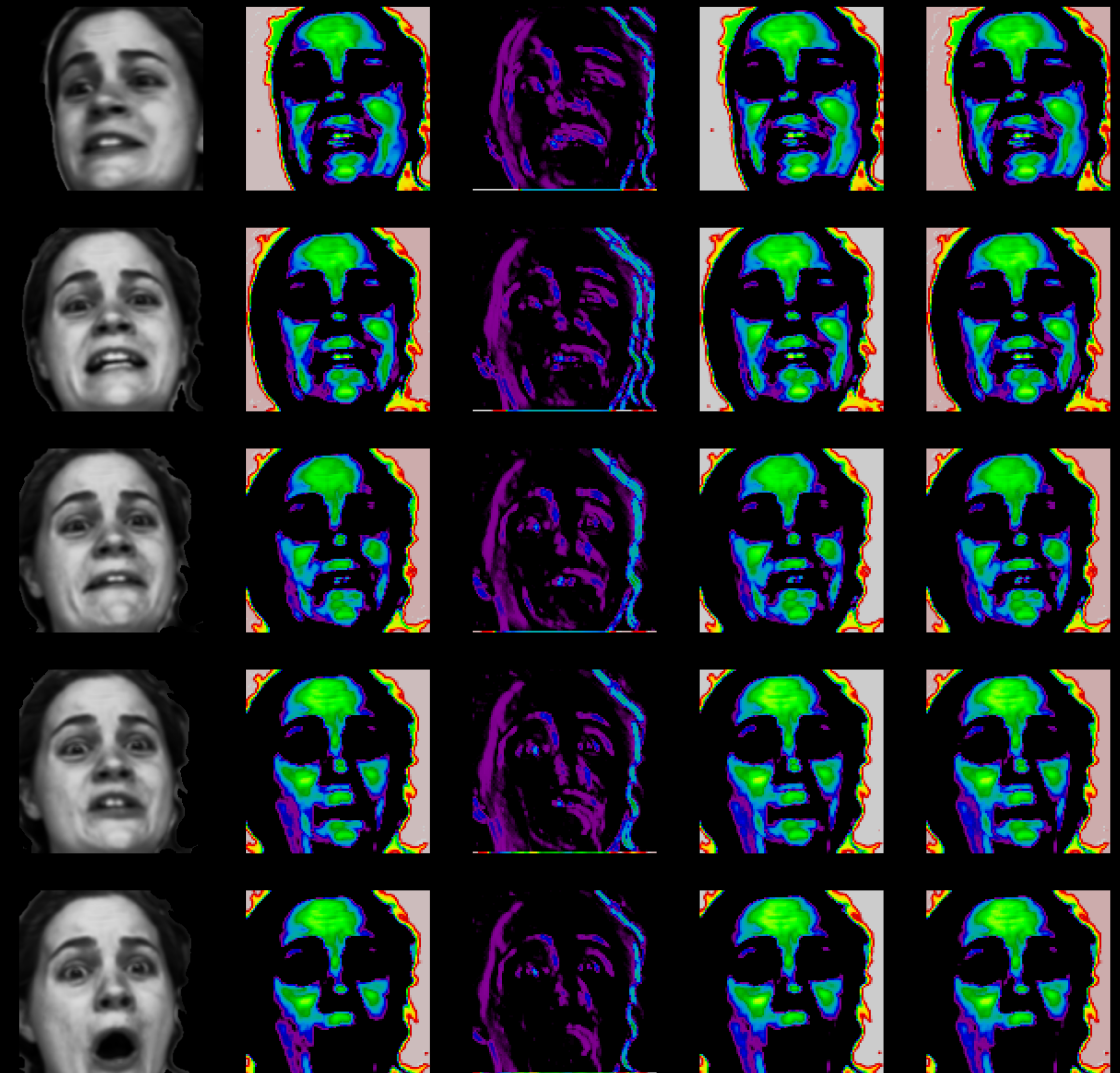
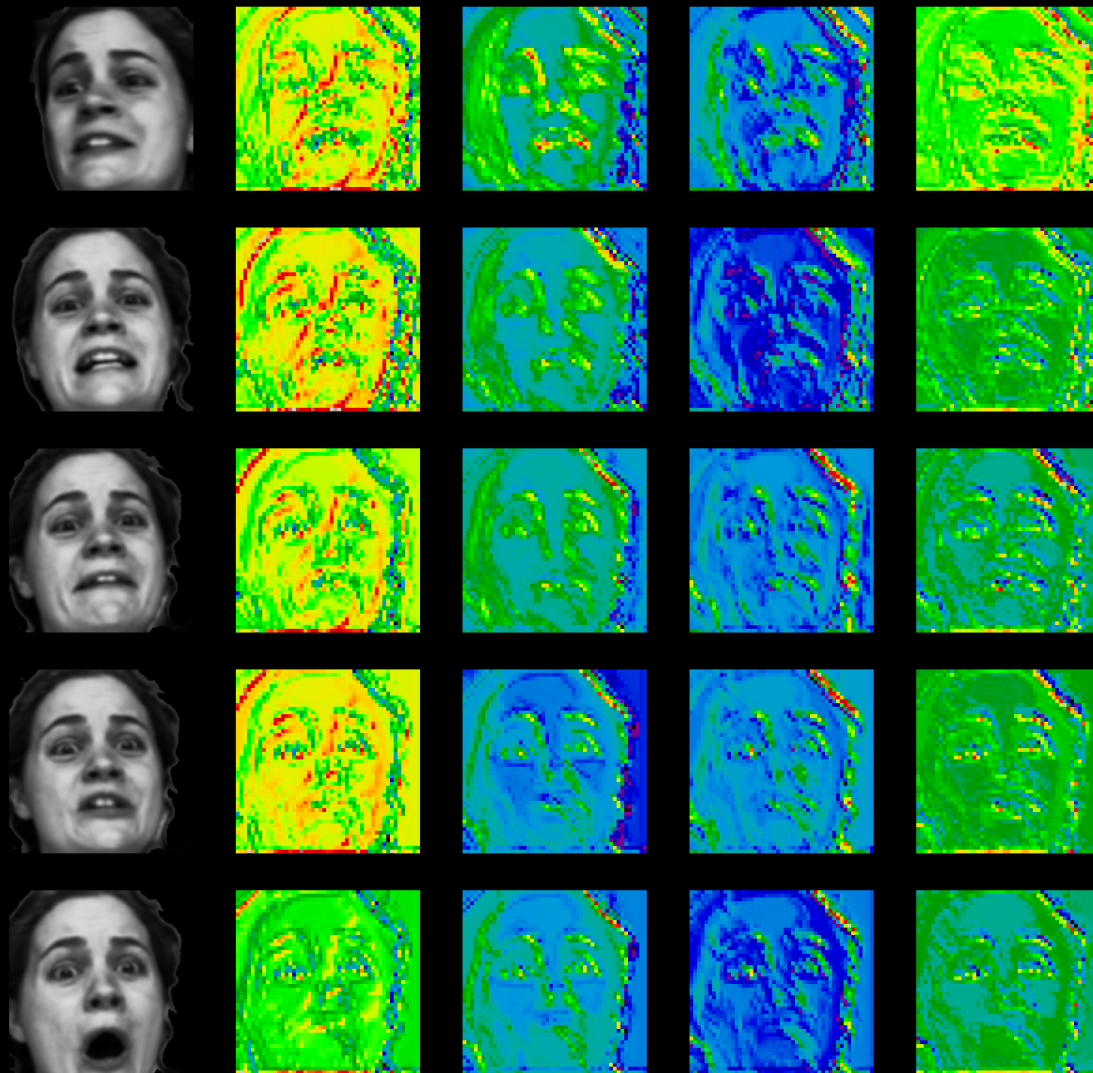


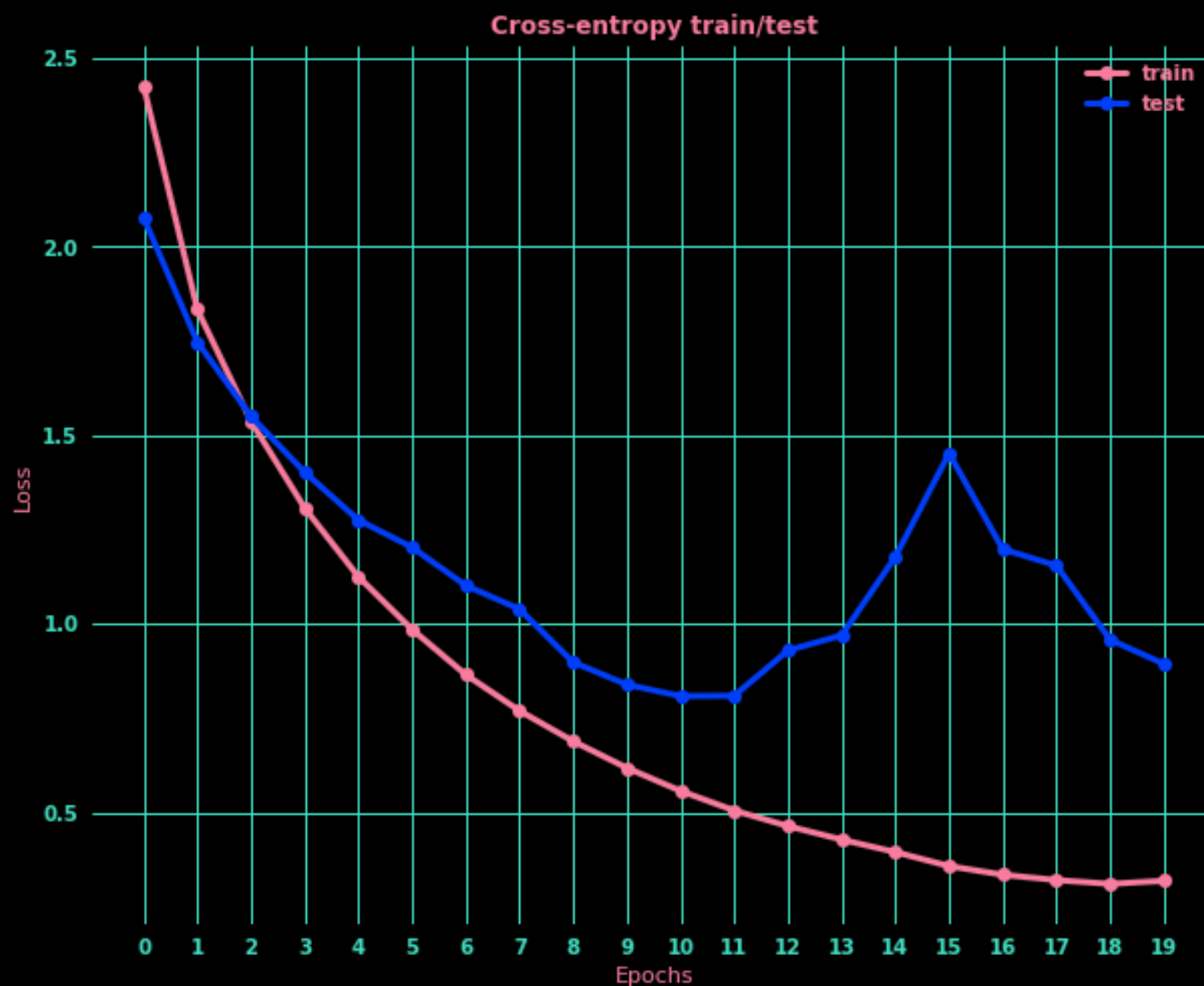
► Гауссове згладжування



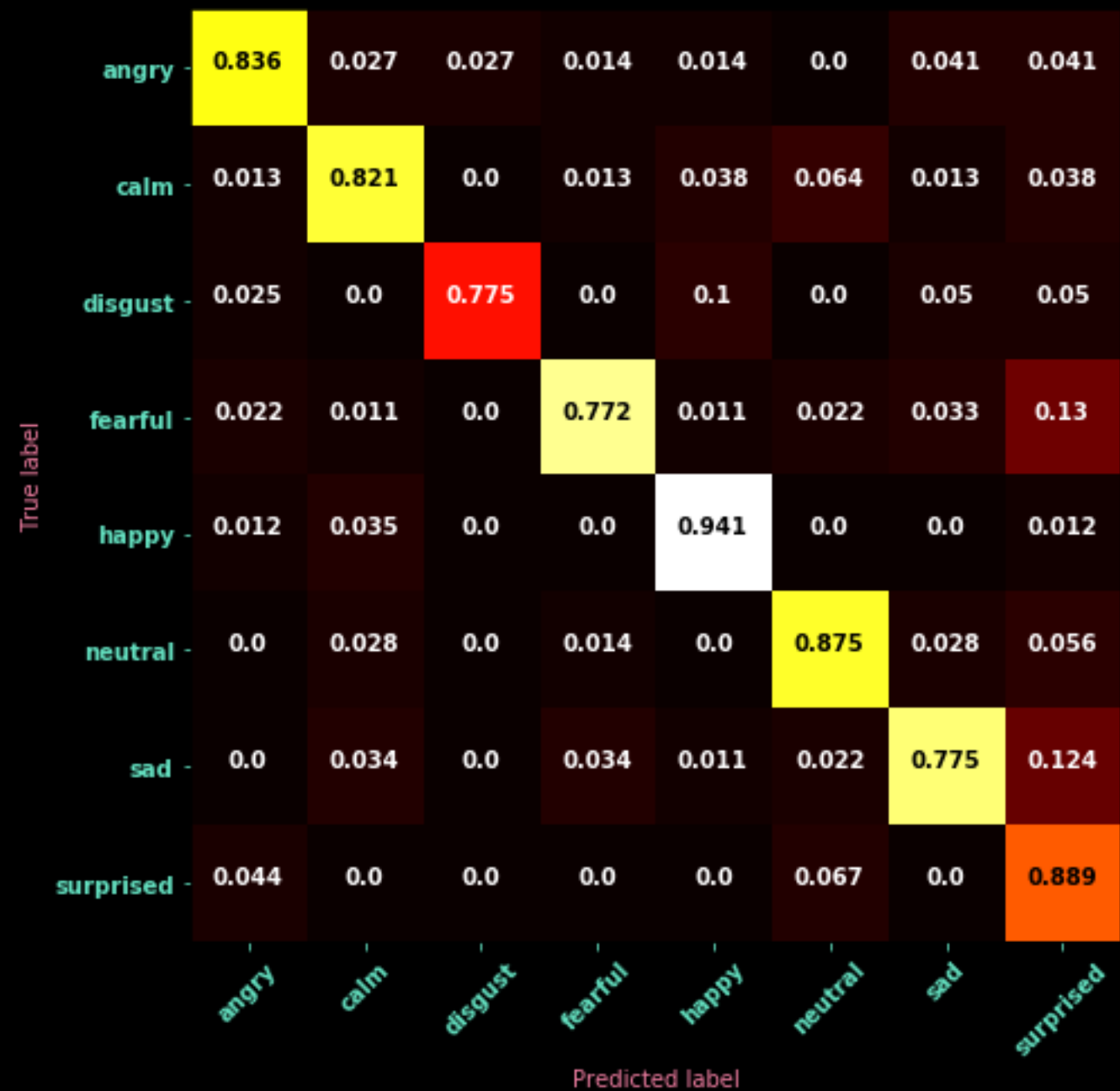


- ▶ Між блоками є також шари нелінійностей (ReLU) та технічні шари (BatchNormalization, Dropout)
- ▶ Замість звичайних блоків використовуються залишкові блоки





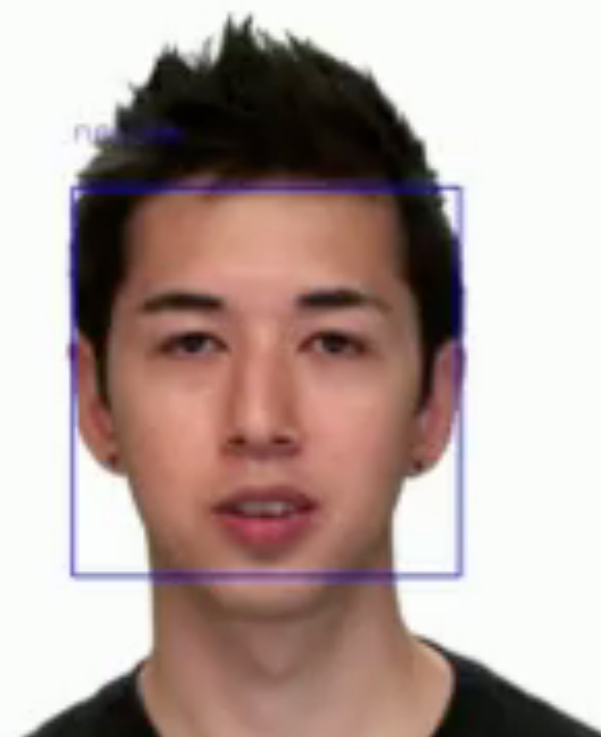
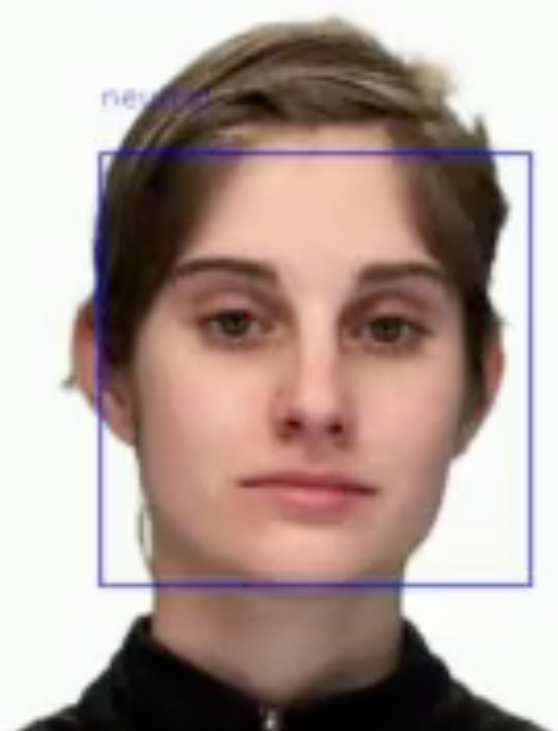
- ▶ Тренувальна вибірка - 0.95
- ▶ Тестова вибірка - 0.8



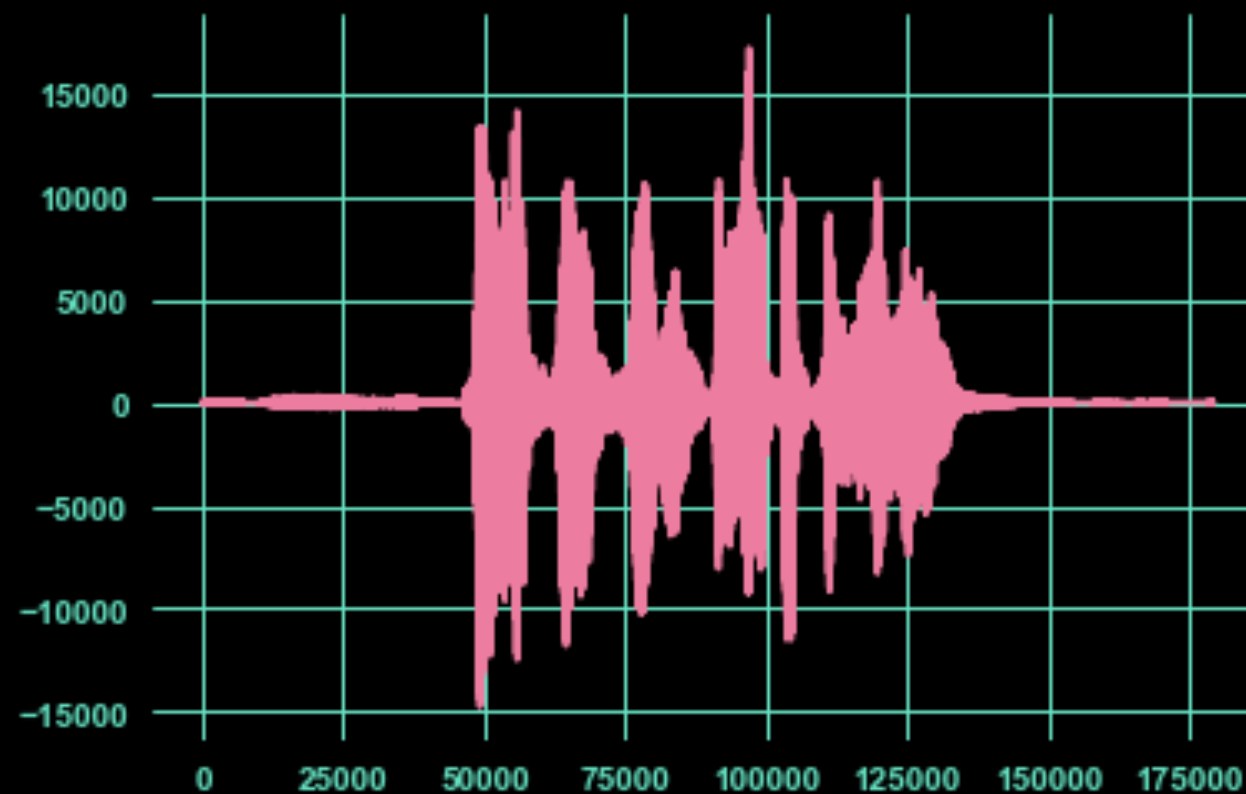
$$C_{ij} = \mathbb{P}[l^* = j \mid l = i]$$

$$C_{ij}^* = \frac{C_{ij}}{\sum_{j=1}^n C_{ij}}$$

$$Accuracy = \frac{tr(C)}{\sum_{i=1}^n \sum_{j=1}^n C_{ij}}$$



- ▶ Просодичні характеристики (енергія, f_0)
- ▶ Спектральні характеристики (MFCC, спектрограми)



- Фреймування вхідного аудіо:

$$\{x_t\}, t = 0, \dots, T$$

- Згладжування фрейму вікном Хеммінга:

$$x_t(n) = x_t(n) * w(n), n = 0, \dots, N - 1$$

$$w(n) = 0.53836 - 0.46164 \cos\left(\frac{2\pi n}{N - 1}\right)$$

- Перетворення Фур'є

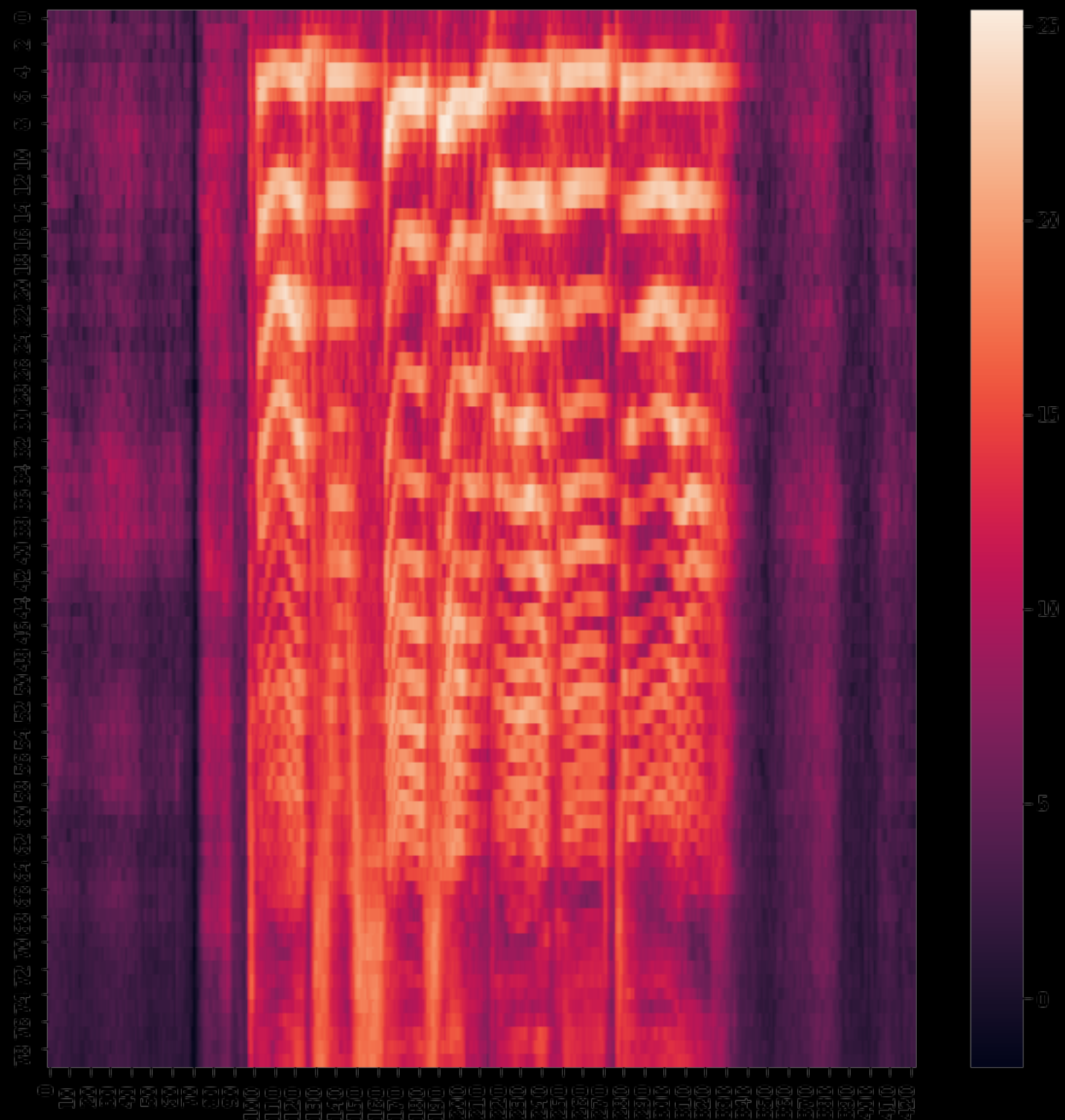
$$X(k) = \sum_{n=0}^{N-1} x_t(n) e^{-\frac{2\pi i}{N-1} nk}, k = 0, \dots, N - 1, t = 0, \dots, T$$

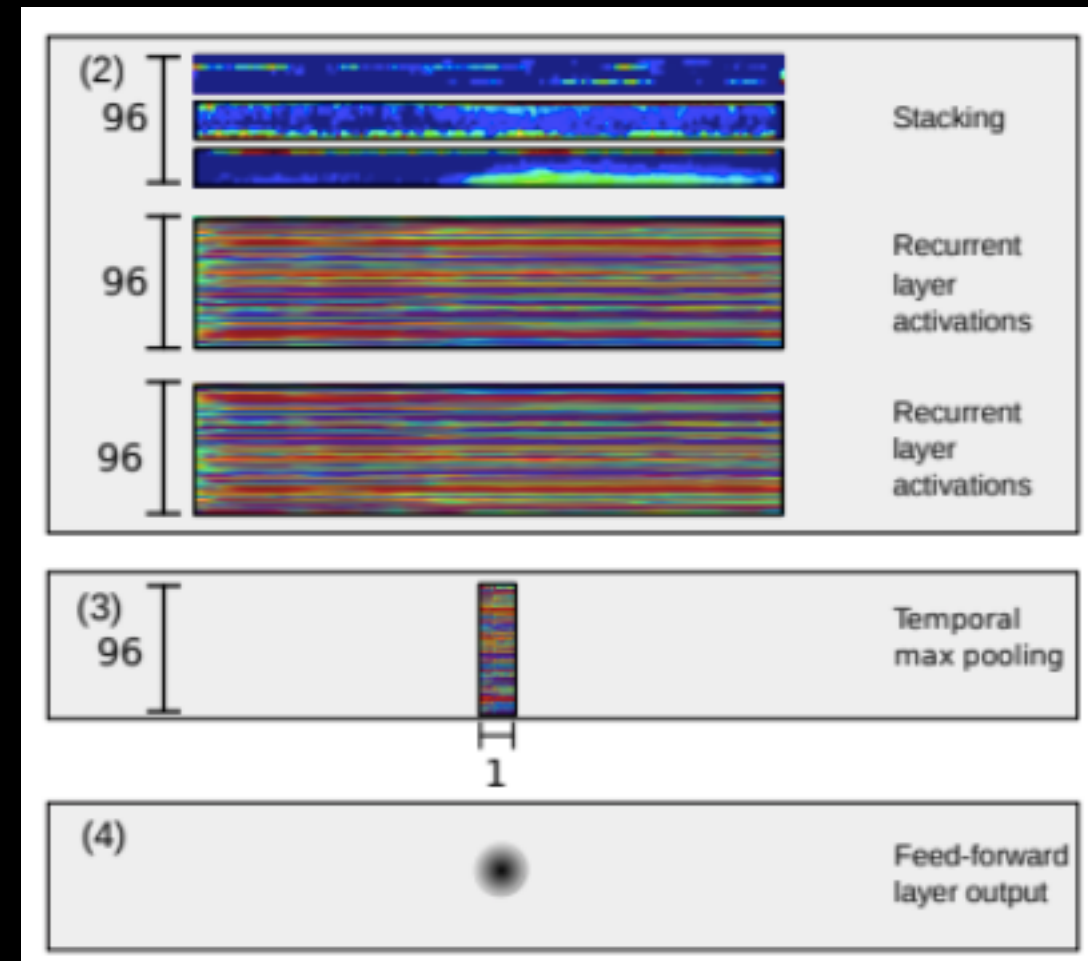
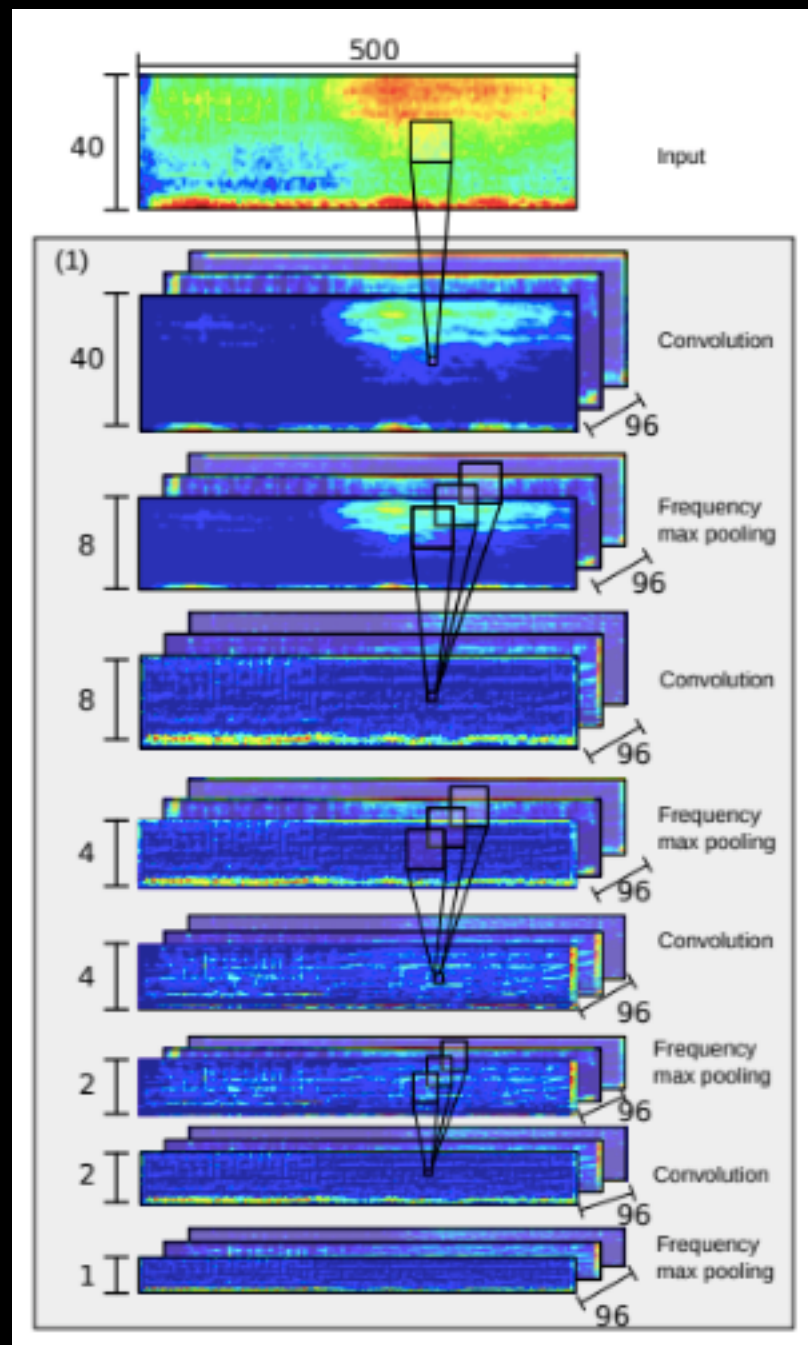
- Створюємо фільтри:

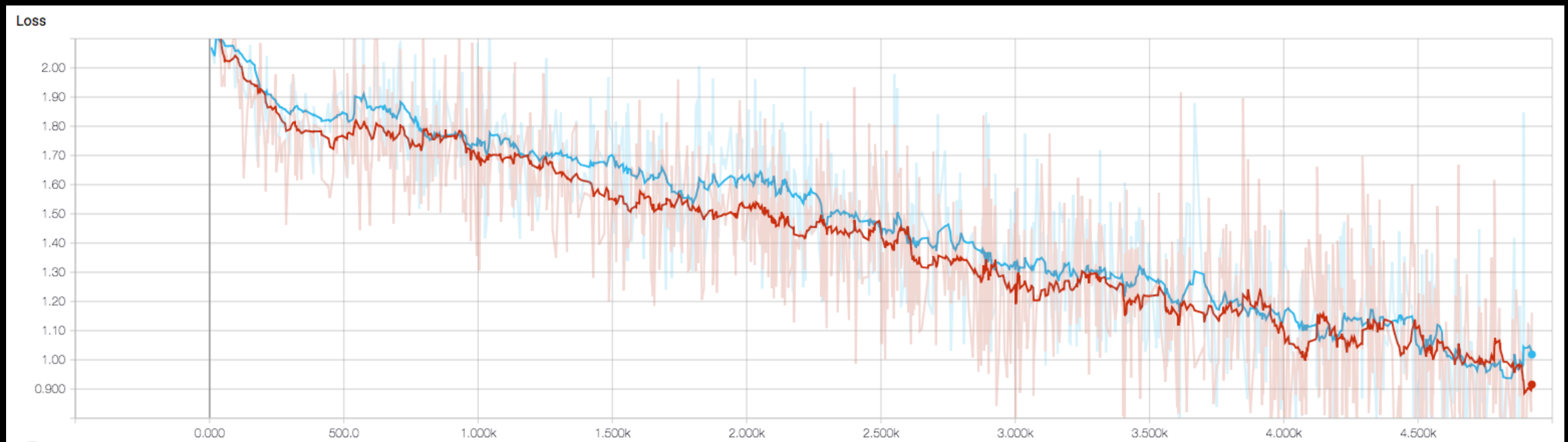
$$H_m = \begin{cases} 0, k < f(m-1) \\ \frac{k - f(m-1)}{f(m) - f(m-1)}, f(m-1) \leq k < f(m) \\ \frac{f(m+1) - k}{f(m+1) - f(m)}, f(m) \leq k \leq f(m+1) \\ 0, k > f(m+1) \end{cases}$$

- Отримуємо енергії для кожного вікна:

$$S_t(m) = \ln\left(\sum_{k=0}^{N-1} |X_t(k)|^2 H_m(k)\right), m = 0, \dots, M$$







- ▶ Тренувальна вибірка - 0.81
- ▶ Тестова вибірка - 0.75



happy.wav

- ▶ Prob 0.00906, angry
- ▶ Prob 0.00505, calm
- ▶ Prob 0.00111, disgust
- ▶ Prob 0.00062, fearful
- ▶ Prob 0.75538, happy
- ▶ Prob 0.20468, neutral
- ▶ Prob 0.02045, sad
- ▶ Prob 0.00364, surprised

sad.wav

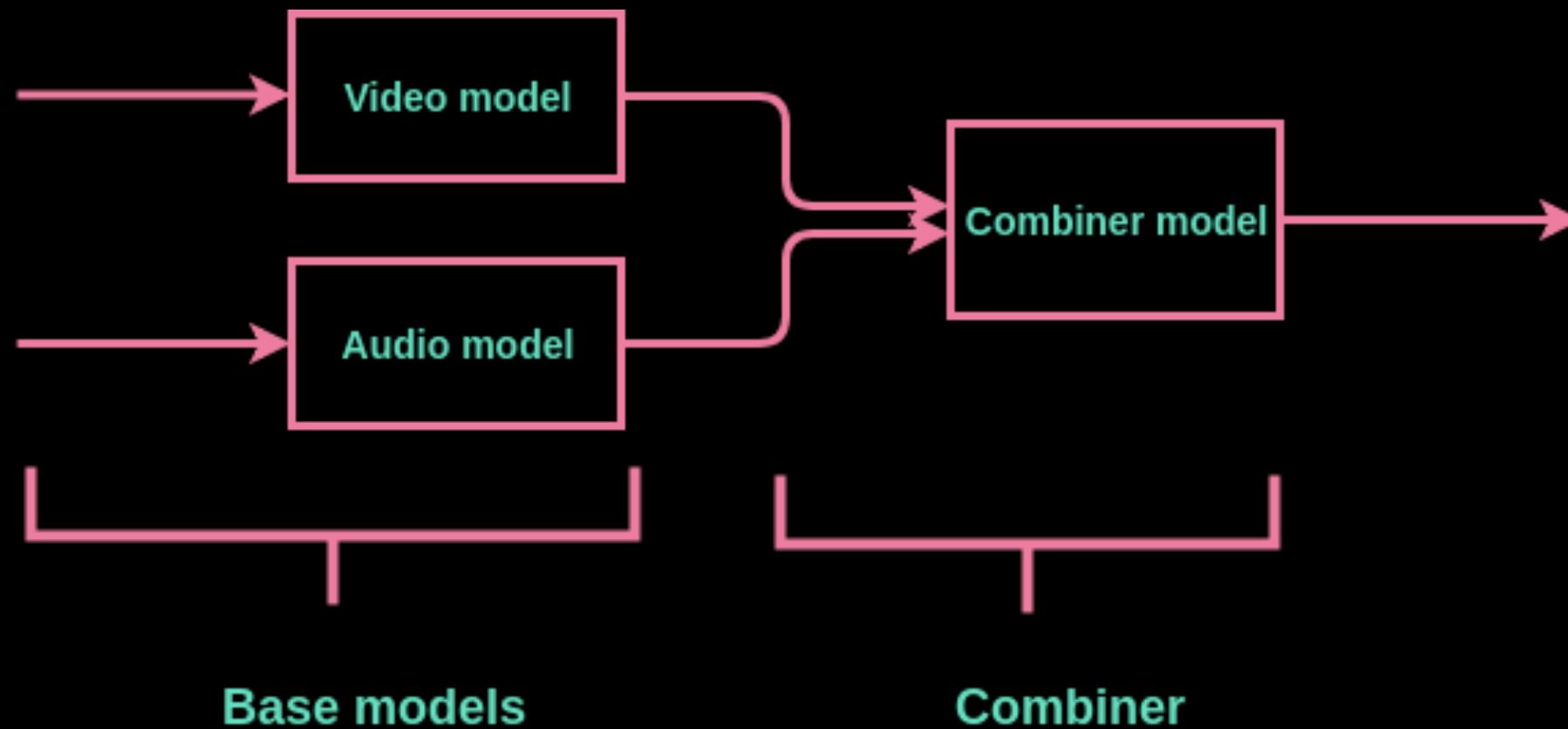
- ▶ Prob 0.00739, angry
- ▶ Prob 0.09296, calm
- ▶ Prob 0.04937, disgust
- ▶ Prob 0.00393, fearful
- ▶ Prob 0.13521, happy
- ▶ Prob 0.11970, neutral
- ▶ Prob 0.48490, sad
- ▶ Prob 0.10654, surprised

fear.wav

- ▶ Prob 0.00908, angry
- ▶ Prob 0.00434, calm
- ▶ Prob 0.00232, disgust
- ▶ Prob 0.80547, fearful
- ▶ Prob 0.11106, happy
- ▶ Prob 0.00208, neutral
- ▶ Prob 0.02724, sad
- ▶ Prob 0.03841, surprised

angry.wav

- ▶ Prob 0.44661, angry
- ▶ Prob 0.00974, calm
- ▶ Prob 0.02614, disgust
- ▶ Prob 0.04990, fearful
- ▶ Prob 0.21963, happy
- ▶ Prob 0.04308, neutral
- ▶ Prob 0.19184, sad
- ▶ Prob 0.01306, surprised





► Тренувальна
вибірка - 0.96

► Тестова вибірка -
0.86

- ▶ Досліджено основні принципи роботи нейронних мереж
- ▶ Розглянуті найефективніші методи глибокого навчання для класифікації відео та аудіо
- ▶ Побудовані нейронні мережі для класифікації емоцій з відео та аудіо
- ▶ Побудований ансамбль моделей, який допоміг значно покращити результат
- ▶ Побудовані моделі адаптовані до зйомки з фронтальної відеокамери та мікрофону ноутбуків

ДЯКУЄМО ЗА УВАГУ
