

ECE 601: Homework 3

Wei Xu

Email: xu1639@purdue.edu

Due date: 11:59 PM, Jan. 29, 2024
(Spring 2024)

1 Becoming Familiar with the Primer

1.1 Comparing the multi-neuron case with the handcrafted network and the `torch.nn` based network

Note: The code used for this subsection is borrowed from the `ComputationalGraphPrimer` package (<https://engineering.purdue.edu/kak/distCGP/>).

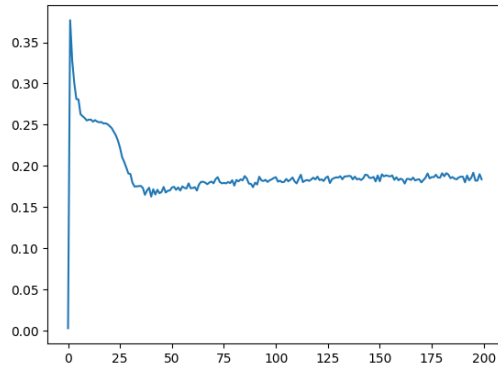
The last ten lines of the loss output from the handcrafted network are shown below. And the loss curve is shown in Fig. 1(a).

```
...
[iter=19001]  loss = 0.1868
[iter=19101]  loss = 0.1800
[iter=19201]  loss = 0.1883
[iter=19301]  loss = 0.1826
[iter=19401]  loss = 0.1855
[iter=19501]  loss = 0.1917
[iter=19601]  loss = 0.1822
[iter=19701]  loss = 0.1822
[iter=19801]  loss = 0.1899
[iter=19901]  loss = 0.1838
```

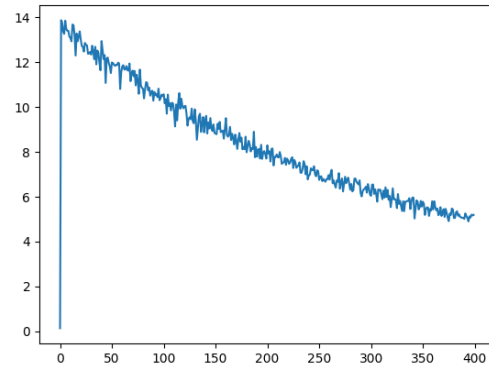
The last ten lines of the loss output from the `torch.nn` network are shown below. And the loss curve is shown in Fig. 1(b).

```
...
[iter=39001]  loss = 5.0162
[iter=39101]  loss = 5.2571
[iter=39201]  loss = 5.1446
[iter=39301]  loss = 5.0601
[iter=39401]  loss = 4.9013
[iter=39501]  loss = 5.1261
[iter=39601]  loss = 5.0779
[iter=39701]  loss = 5.2043
[iter=39801]  loss = 5.1673
[iter=39901]  loss = 5.1952
```

It is noticed from Fig. 1, the `torch.nn` network converges more stable and the loss decreases all the way, but the handcrafted network shows some unstable behaviors, like the loss increasing during the late training.



(a) handcrafted network



(b) `torch.nn` network

Figure 1: Loss curves for multi-neuron case.

1.2 Comparing the one-neuron case with the handcrafted network and the `torch.nn` based network

Note: The code used for this subsection is borrowed from the `ComputationalGraphPrimer` package (<https://engineering.purdue.edu/kak/distCGP/>).

The last ten lines of the loss output from the handcrafted network are shown below. And the loss curve is shown in Fig. 2(a).

```
...
[iter=39001]  loss = 0.1884
[iter=39101]  loss = 0.1906
[iter=39201]  loss = 0.1924
[iter=39301]  loss = 0.1929
[iter=39401]  loss = 0.1960
[iter=39501]  loss = 0.1883
[iter=39601]  loss = 0.1885
[iter=39701]  loss = 0.1932
[iter=39801]  loss = 0.1877
[iter=39901]  loss = 0.1906
```

To test the `torch.nn` network, some modifications are needed in `verify_with_torchnn.py` to run for this one-neuron case. The modifications and the last ten lines of the loss output are shown below. And the loss curve is shown in Fig. 1(b).

```
...
cgp = ComputationalGraphPrimer(
    ...
    # learning_rate = 1e-6, # For the multi-neuron option
    # ↪ below
    ...
    learning_rate = 5 * 1e-2, # Also for the one-neuron
    # ↪ option below
    ...
)
...
```

```

cgp.run_training_with_torchnn('one_neuron', training_data) ## (A)
    ↪ REMEMBER to also change learning_rate above
# cgp.run_training_with_torchnn('multi_neuron', training_data) ## (B)
    ↪ REMEMBER to also change learning_rate above

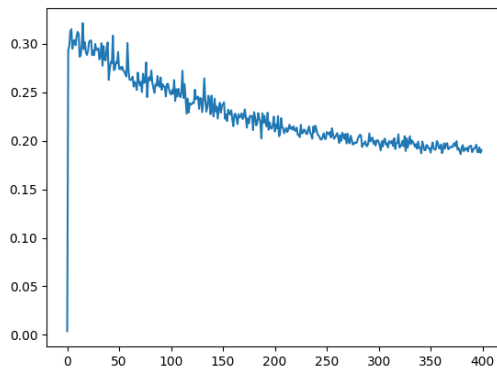
```

```

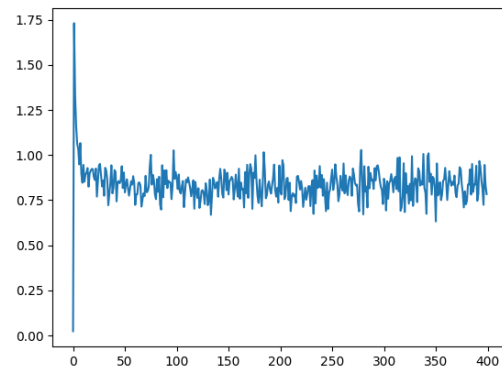
...
[iter=39001]    loss = 0.7465
[iter=39101]    loss = 0.7844
[iter=39201]    loss = 0.9668
[iter=39301]    loss = 0.9348
[iter=39401]    loss = 0.8580
[iter=39501]    loss = 0.8124
[iter=39601]    loss = 0.7243
[iter=39701]    loss = 0.9452
[iter=39801]    loss = 0.8305
[iter=39901]    loss = 0.7839

```

It is noticed from Fig. 2, the `torch.nn` network converges much faster, but the handcrafted network converges slower.



(a) handcrafted network



(b) `torch.nn` network

Figure 2: Loss curves for one-neuron case.

1.3 Understanding of SGD+ and Adam

SGD with momentum (SGD+). According to [Slides 108 and 109 in Week 3](#), using the Vanilla Gradient Descent is prone to oscillations in the convergence curve when encountering "valleys". One solution to this problem is to apply momentum, namely SGD with momentum (SGD+). Specifically, the idea is to add momentum to the gradient descent process, making the updating faster in the dimension in which the gradient is unchanged and the updating slower in the dimension in which the gradient has changed. This speeds up convergence and reduces oscillations. The update process is

$$v_t = \mu * v_{t-1} + g_t \quad (1)$$

$$p_t = p_{t-1} - \text{lr} * v_t \quad (2)$$

where v_t is the step size at t (typically $v_0 = 0$), g_t is the gradient at t , p_t is the model weights at t , lr is the learning rate, and $\mu \in [0, 1]$ is the momentum factor. μ controls how much the previous step

size contributes to the current step size. A larger μ means that the previous step size is more heavily weighted in the current step size calculation. The current direction of descent is determined by the combination of the direction of descent accumulated from the history and the current gradient. And the closer to the current moment, the more influence it exerts. The advantages of SGD+ are the ability to pass local minima, the ability to speed up convergence, as well as the ability to suppress oscillations during gradient descent. When checking the PyTorch function of SGD (<https://pytorch.org/docs/stable/generated/torch.optim.SGD.html>), this parameter is set as 0 by default. But I think a value of 0.5 is good to start from.

Adaptive Moment Estimation (Adam). SGD+ considers the first-order momentum, to which Adam adds the second-order momentum. [Slides 117 and 118 in Week 3](#) provides good explanation to this optimizer. Firstly, calculate the biased first moment estimate m_t with

$$m_t = \beta_1 * m_{t-1} + (1 - \beta_1) * g_t \quad (3)$$

where β_1 is the first exponential decay rate, controlling computing running averages of the gradient. Then, calculate the biased second raw moment estimate v_t with

$$v_t = \beta_2 * v_{t-1} + (1 - \beta_2) * (g_t)^2 \quad (4)$$

where β_2 is the second exponential decay rate, controlling computing running averages of the squared gradient. Notice that the initial values $m_0 = 0$ and $v_0 = 0$, then the following values m_t and v_t ($t \leq 1$) are biased toward 0, especially during the beginning training phase. So these two values are corrected by

$$\hat{m}_t = \frac{m_t}{1 - \beta_1^t} \quad (5)$$

$$\hat{v}_t = \frac{v_t}{1 - \beta_2^t} \quad (6)$$

Finally, the model weights are updated by

$$p_t = p_{t-1} - \text{lr} * \frac{\hat{m}_t}{\sqrt{\hat{v}_t + \epsilon}} \quad (7)$$

where ϵ is incorporated to avoid the denominator being 0, which is set as 10^{-8} in PyTorch by default. With the second moment item, Adam is more robust and can help converge faster compared with SGD+.

2 Programming Task

2.1 One-neuron classifier

Note: The code for this subsection is borrowed from the [ComputationalGraphPrimer](https://engineering.purdue.edu/kak/distCGP/) package (<https://engineering.purdue.edu/kak/distCGP/>).

The code for this task is shown below. To make the code concise, I only put the comments in the lines that I modified. The `OneNeuronSGDCGP` class is for Vanilla Gradient Descent. The `run_training_loop_one_neuron_model` method is rewritten to return the loss record. The `OneNeuronSGDPlusCGP` class is for SGD+. A new parameter, momentum factor, is added in the `__init__` method. The `run_training_loop_one_neuron_model` method is rewritten to return the loss record. The `backprop_and_update_params_one_neuron_model` method is rewritten to apply

the SGD+ weight update equations. The `OneNeuronAdamCGP` class is for Adam. Three new parameters, first exponential decay rate, second exponential decay rate, and epsilon, are added in the `__init__` method. The `run_training_loop_one_neuron_model` method is rewritten to return the loss record. The `backprop_and_update_params_one_neuron_model` method is rewritten to apply the Adam weight update equations.

```
...
class OneNeuronSGDCGP(ComputationalGraphPrimer):
    def __init__(self, *args, **kwargs):
        super().__init__(*args, **kwargs) ## inherit from CGP
    def run_training_loop_one_neuron_model(self, training_data):
        self.vals_for_learnable_params = {param: random.uniform(0,1)
        ↪ for param in self.learnable_params}
        self.bias = random.uniform(0,1)
        class DataLoader:
            def __init__(self, training_data, batch_size):
                self.training_data = training_data
                self.batch_size = batch_size
                self.class_0_samples = [(item, 0) for item in self.
                ↪ training_data[0]]
                self.class_1_samples = [(item, 1) for item in self.
                ↪ training_data[1]]
            def __len__(self):
                return len(self.training_data[0]) + len(self.
                ↪ training_data[1])
            def _getitem(self):
                cointoss = random.choice([0,1])
                if cointoss == 0:
                    return random.choice(self.class_0_samples)
                else:
                    return random.choice(self.class_1_samples)
            def getbatch(self):
                batch_data, batch_labels = [], []
                maxval = 0.0
                for _ in range(self.batch_size):
                    item = self._getitem()
                    if np.max(item[0]) > maxval:
                        maxval = np.max(item[0])
                    batch_data.append(item[0])
                    batch_labels.append(item[1])
                batch_data = [item/maxval for item in batch_data]
                batch = [batch_data, batch_labels]
                return batch
        data_loader = DataLoader(training_data, batch_size=self.
        ↪ batch_size)
        loss_running_record = []
        i = 0
        avg_loss_over_iterations = 0.0
        for i in range(self.training_iterations):
            data = data_loader.getbatch()
            data_tuples_in_batch = data[0]
            class_labels_in_batch = data[1]
            y_preds, deriv_sigmoids = self.
            ↪ forward_prop_one_neuron_model(data_tuples_in_batch)
```

```

        loss = sum([(abs(class_labels_in_batch[i] - y_preds[i]))**2
        ↪ for i in range(len(class_labels_in_batch))])
    avg_loss_over_iterations += loss / float(len(
    ↪ class_labels_in_batch))
    if i%(self.display_loss_how_often) == 0:
        avg_loss_over_iterations /= self.display_loss_how_often
        loss_running_record.append(avg_loss_over_iterations)
        print("[iter=%d] loss = %.4f" % (i+1,
        ↪ avg_loss_over_iterations))
        avg_loss_over_iterations = 0.0
    y_errors_in_batch = list(map(operator.sub,
    ↪ class_labels_in_batch, y_preds))
    self.backprop_and_update_params_one_neuron_model(
    ↪ data_tuples_in_batch, y_preds, y_errors_in_batch,
    ↪ deriv_sigmoids)
    ## removed plt functions
    return loss_running_record ## return the loss record

class OneNeuronSGDPlusCGP(ComputationalGraphPrimer):
    def __init__(self, mu, *args, **kwargs):
        super().__init__(*args, **kwargs) ## inherit from CGP
        self.mu = mu ## momentum factor
    def run_training_loop_one_neuron_model(self, training_data):
        self.vals_for_learnable_params = {param: random.uniform(0,1)
        ↪ for param in self.learnable_params}
        self.bias = random.uniform(0,1)
        self.previous_v_vals = {param: 0 for param in self.
        ↪ learnable_params} ## initialize v_0 for weights
        self.previous_v_bias = 0 ## initialize v_0 for bias
    class DataLoader:
        def __init__(self, training_data, batch_size):
            self.training_data = training_data
            self.batch_size = batch_size
            self.class_0_samples = [(item, 0) for item in self.
            ↪ training_data[0]]
            self.class_1_samples = [(item, 1) for item in self.
            ↪ training_data[1]]
        def __len__(self):
            return len(self.training_data[0]) + len(self.
            ↪ training_data[1])
        def _getitem(self):
            cointoss = random.choice([0,1])
            if cointoss == 0:
                return random.choice(self.class_0_samples)
            else:
                return random.choice(self.class_1_samples)
        def getbatch(self):
            batch_data, batch_labels = [], []
            maxval = 0.0
            for _ in range(self.batch_size):
                item = self._getitem()
                if np.max(item[0]) > maxval:
                    maxval = np.max(item[0])
            batch_data.append(item[0])

```

```

        batch_labels.append(item[1])
        batch_data = [item/maxval for item in batch_data]
        batch = [batch_data, batch_labels]
        return batch
data_loader = DataLoader(training_data, batch_size=self.
    ↪ batch_size)
loss_running_record = []
i = 0
avg_loss_over_iterations = 0.0
for i in range(self.training_iterations):
    data = data_loader.getbatch()
    data_tuples_in_batch = data[0]
    class_labels_in_batch = data[1]
    y_preds, deriv_sigmoids = self.
        ↪ forward_prop_one_neuron_model(data_tuples_in_batch)
    loss = sum([(abs(class_labels_in_batch[i] - y_preds[i]))**2
        ↪ for i in range(len(class_labels_in_batch))])
    avg_loss_over_iterations += loss / float(len(
        ↪ class_labels_in_batch))
    if i%(self.display_loss_how_often) == 0:
        avg_loss_over_iterations /= self.display_loss_how_often
        loss_running_record.append(avg_loss_over_iterations)
        print("[iter=%d] loss = %.4f" % (i+1,
            ↪ avg_loss_over_iterations))
        avg_loss_over_iterations = 0.0
    y_errors_in_batch = list(map(operator.sub,
        ↪ class_labels_in_batch, y_preds))
    self.backprop_and_update_params_one_neuron_model(
        ↪ data_tuples_in_batch, y_preds, y_errors_in_batch,
        ↪ deriv_sigmoids)
    ## removed plt functions
    return loss_running_record ## return the loss record
def backprop_and_update_params_one_neuron_model(self,
    ↪ data_tuples_in_batch, predictions, y_errors_in_batch,
    ↪ deriv_sigmoids):
    input_vars = self.independent_vars
    input_vars_to_param_map = self.var_to_var_param[self.
        ↪ output_vars[0]]
    param_to_vars_map = {param : var for var, param in
        ↪ input_vars_to_param_map.items()}
    vals_for_learnable_params = self.vals_for_learnable_params
    for i,param in enumerate(self.vals_for_learnable_params):
        partial_of_loss_wrt_param = 0.0
        for j in range(self.batch_size):
            vals_for_input_vars_dict = dict(zip(input_vars, list(
                ↪ data_tuples_in_batch[j])))
            partial_of_loss_wrt_param += - y_errors_in_batch[j] *
                ↪ vals_for_input_vars_dict[param_to_vars_map[param]
                ↪ ]] * deriv_sigmoids[j]
        partial_of_loss_wrt_param /= float(self.batch_size) ##
            ↪ average after the batch finishing
        self.previous_v_vals[param] = self.mu * self.
            ↪ previous_v_vals[param] + partial_of_loss_wrt_param ##
            ↪ update momentum memory

```

```

        step = self.learning_rate * self.previous_v_vals[param] ##
        ↪ scaled step size
        self.vals_for_learnable_params[param] -= step
    partial_of_loss = 0 ## initialize loss count
    for j in range(self.batch_size): ## loop for a batch
        partial_of_loss -= y_errors_in_batch[j] * deriv_sigmoids[j]
    partial_of_loss /= float(self.batch_size) ## average after the
    ↪ batch finishing
    self.previous_v_bias = self.mu * self.previous_v_bias +
    ↪ partial_of_loss ## update momentum memory
    self.bias -= self.learning_rate * self.previous_v_bias ##
    ↪ Update the bias

class OneNeuronAdamCGP(ComputationalGraphPrimer):
    def __init__(self, beta_1, beta_2, epsilon, *args, **kwargs):
        super().__init__(*args, **kwargs) ## inherit from CGP
        self.beta_1 = beta_1 ## first exponential decay rate
        self.beta_2 = beta_2 ## second exponential decay rate
        self.epsilon = epsilon ## epsilon
    def run_training_loop_one_neuron_model(self, training_data):
        self.vals_for_learnable_params = {param: random.uniform(0,1)
        ↪ for param in self.learnable_params}
        self.bias = random.uniform(0,1)
        self.previous_m_vals = {param: 0 for param in self.
        ↪ learnable_params} ## initialize m_0 for weights
        self.previous_m_bias = 0 ## initialize m_0 for bias
        self.previous_v_vals = {param: 0 for param in self.
        ↪ learnable_params} ## initialize m_0 for weights
        self.previous_v_bias = 0 ## initialize m_0 for bias
    class DataLoader:
        def __init__(self, training_data, batch_size):
            self.training_data = training_data
            self.batch_size = batch_size
            self.class_0_samples = [(item, 0) for item in self.
            ↪ training_data[0]]
            self.class_1_samples = [(item, 1) for item in self.
            ↪ training_data[1]]
        def __len__(self):
            return len(self.training_data[0]) + len(self.
            ↪ training_data[1])
        def _getitem(self):
            cointoss = random.choice([0,1])
            if cointoss == 0:
                return random.choice(self.class_0_samples)
            else:
                return random.choice(self.class_1_samples)
        def getbatch(self):
            batch_data, batch_labels = [], []
            maxval = 0.0
            for _ in range(self.batch_size):
                item = self._getitem()
                if np.max(item[0]) > maxval:
                    maxval = np.max(item[0])
            batch_data.append(item[0])

```



```

        batch_labels.append(item[1])
        batch_data = [item/maxval for item in batch_data]
        batch = [batch_data, batch_labels]
        return batch
data_loader = DataLoader(training_data, batch_size=self.
    ↪ batch_size)
loss_running_record = []
i = 0
avg_loss_over_iterations = 0.0
for i in range(self.training_iterations):
    data = data_loader.getbatch()
    data_tuples_in_batch = data[0]
    class_labels_in_batch = data[1]
    y_preds, deriv_sigmoids = self.
        ↪ forward_prop_one_neuron_model(data_tuples_in_batch)
    loss = sum([(abs(class_labels_in_batch[i] - y_preds[i]))**2
        ↪ for i in range(len(class_labels_in_batch))])
    avg_loss_over_iterations += loss / float(len(
        ↪ class_labels_in_batch))
    if i%(self.display_loss_how_often) == 0:
        avg_loss_over_iterations /= self.display_loss_how_often
        loss_running_record.append(avg_loss_over_iterations)
        print("[iter=%d] loss = %.4f" % (i+1,
            ↪ avg_loss_over_iterations))
        avg_loss_over_iterations = 0.0
    y_errors_in_batch = list(map(operator.sub,
        ↪ class_labels_in_batch, y_preds))
    self.backprop_and_update_params_one_neuron_model(
        ↪ data_tuples_in_batch, y_preds, y_errors_in_batch,
        ↪ deriv_sigmoids, i+1)
    ## removed plt functions
    return loss_running_record ## return the loss record
def backprop_and_update_params_one_neuron_model(self,
    ↪ data_tuples_in_batch, predictions, y_errors_in_batch,
    ↪ deriv_sigmoids, iter): ## add iteration index as aparameter
    input_vars = self.independent_vars
    input_vars_to_param_map = self.var_to_var_param[self.
        ↪ output_vars[0]]
    param_to_vars_map = {param : var for var, param in
        ↪ input_vars_to_param_map.items()}
    vals_for_learnable_params = self.vals_for_learnable_params
    for i,param in enumerate(self.vals_for_learnable_params):
        partial_of_loss_wrt_param = 0.0
        for j in range(self.batch_size):
            vals_for_input_vars_dict = dict(zip(input_vars, list(
                ↪ data_tuples_in_batch[j])))
            partial_of_loss_wrt_param += - y_errors_in_batch[j] *
                ↪ vals_for_input_vars_dict[param_to_vars_map[param]
                ↪ ]] * deriv_sigmoids[j]
        partial_of_loss_wrt_param /= float(self.batch_size) ##
            ↪ average after the batch finishing
    self.previous_m_vals[param] = self.beta_1 * self.
        ↪ previous_m_vals[param] + (1 - self.beta_1) *
        ↪ partial_of_loss_wrt_param ## update first momentum

```

```

        ↪ memory
        self.previous_v_vals[param] = self.beta_2 * self.
        ↪ previous_v_vals[param] + (1 - self.beta_2) *
        ↪ partial_of_loss_wrt_param**2 ## update second
        ↪ momentum memory
        m_t_hat = self.previous_m_vals[param] / (1 - self.beta_1**
        ↪ iter) ## scaled first momentum memory
        v_t_hat = self.previous_v_vals[param] / (1 - self.beta_2**
        ↪ iter) ## scaled second momentum memory
        step = self.learning_rate * m_t_hat / np.sqrt(v_t_hat +
        ↪ self.epsilon) ## scaled step size
        self.vals_for_learnable_params[param] -= step
        partial_of_loss = 0 ## initialize loss count
        for j in range(self.batch_size): ## loop for a batch
            partial_of_loss -= y_errors_in_batch[j] * deriv_sigmoids[j]
        partial_of_loss /= float(self.batch_size) ## average after the
        ↪ batch finishing
        self.previous_m_bias = self.beta_1 * self.previous_m_bias + (1
        ↪ - self.beta_1) * partial_of_loss ## update first momentum
        ↪ memory
        self.previous_v_bias = self.beta_2 * self.previous_v_bias + (1
        ↪ - self.beta_2) * partial_of_loss**2 ## update second
        ↪ momentum memory
        m_t_hat = self.previous_m_bias / (1 - self.beta_1**iter) ##
        ↪ scaled first momentum memory
        v_t_hat = self.previous_v_bias / (1 - self.beta_2**iter) ##
        ↪ scaled second momentum memory
        self.bias -= self.learning_rate * m_t_hat / np.sqrt(v_t_hat +
        ↪ self.epsilon) ## Update the bias
    ...
## plot the learning curves
def plot_learning_curve(loss_SGD, loss_SGDPlus, loss_Adam, title, lr):
    plt.figure(figsize=(6, 4))
    plt.plot(loss_SGD, label='SGD')
    plt.plot(loss_SGDPlus, label='SGD+')
    plt.plot(loss_Adam, label='Adam')
    plt.xlabel('Iterations')
    plt.ylabel('losses')
    plt.legend()
    cf = plt.gcf()
    cf.savefig('%s_%s.jpg'%(title,lr), bbox_inches='tight', \
               format='jpg', dpi=1000)

## run one-neuron case with different optimizer
def run_one_neuron_diff_optimizer(lr):
    cgp_SGD = OneNeuronSGDCGP(
        one_neuron_model = True,
        expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
        output_vars = ['xw'],
        dataset_size = 5000,
        learning_rate = lr,
        training_iterations = 40000,
        batch_size = 8,
        display_loss_how_often = 100,

```

```

        debug = True,
    )
    cgp_SGD.parse_expressions()
    cgp_SGD.display_one_neuron_network()

    cgp_SGDPlus = OneNeuronSGDPlusCGP(
        mu = 0.5,
        one_neuron_model = True,
        expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
        output_vars = ['xw'],
        dataset_size = 5000,
        learning_rate = lr,
        training_iterations = 40000,
        batch_size = 8,
        display_loss_how_often = 100,
        debug = True,
    )
    cgp_SGDPlus.parse_expressions()
    cgp_SGDPlus.display_one_neuron_network()

    cgp_Adam = OneNeuronAdamCGP(
        beta_1 = 0.9,
        beta_2 = 0.99,
        epsilon = 1e-8,
        one_neuron_model = True,
        expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
        output_vars = ['xw'],
        dataset_size = 5000,
        learning_rate = lr,
        training_iterations = 40000,
        batch_size = 8,
        display_loss_how_often = 100,
        debug = True,
    )
    cgp_Adam.parse_expressions()
    cgp_Adam.display_one_neuron_network()

    training_data = cgp_SGD.gen_training_data()
    loss_SGD = cgp_SGD.run_training_loop_one_neuron_model(training_data
        ↪ )
    loss_SGDPlus = cgp_SGDPlus.run_training_loop_one_neuron_model(
        ↪ training_data)
    loss_Adam = cgp_Adam.run_training_loop_one_neuron_model(
        ↪ training_data)

    plot_learning_curve(loss_SGD, loss_SGDPlus, loss_Adam, 'one_neuron'
        ↪ , str(lr))
...
if __name__ == '__main__':
...
    if args.task == '1':
        for lr in [1e-2, 1e-3, 1e-4]:
            run_one_neuron_diff_optimizer(lr)
...

```

Fig. 3 shows the loss curves with the learning rates of 0.01, 0.001, and 0.0001 respectively. Generally speaking, Adam performs the best, and SGD does the worst. When the learning rate is large (as in Fig. 3(a)), Adam converges fastest but its loss curve is not very smooth, which is because it achieves the global minimum point and just oscillates around that point; SGD+ converges slower (but finally it achieves a similar level with Adam) and has more smooth loss curve; SGD diverges eventually, so the smoothness of its loss curve is meaningless. When the learning rate is moderate (as in Fig. 3(b)), Adam still converges fastest and gets the global minimum point; SGD+ and SGD converge slower and do not completely achieve the global minimum point eventually; SGD has the most oscillating loss curve in this case. When the learning rate is small (as in Fig. 3(c)), both Adam and SGD+ converge, but Adam is still faster and has the most smooth loss curve, and SGD even does not converge. Compared with Fig. 2(b), it can be found that `torch.nn` beats these three self-implemented optimizers in terms of the converging speed. The reason is that `PyTorch` is a very well-optimized framework, meaning there may be better optimization mechanisms in backpropagation, such as gradient clipping. In addition, this is almost the simplest neural network, the advantages of Adam and SGD+ can not be performed well.

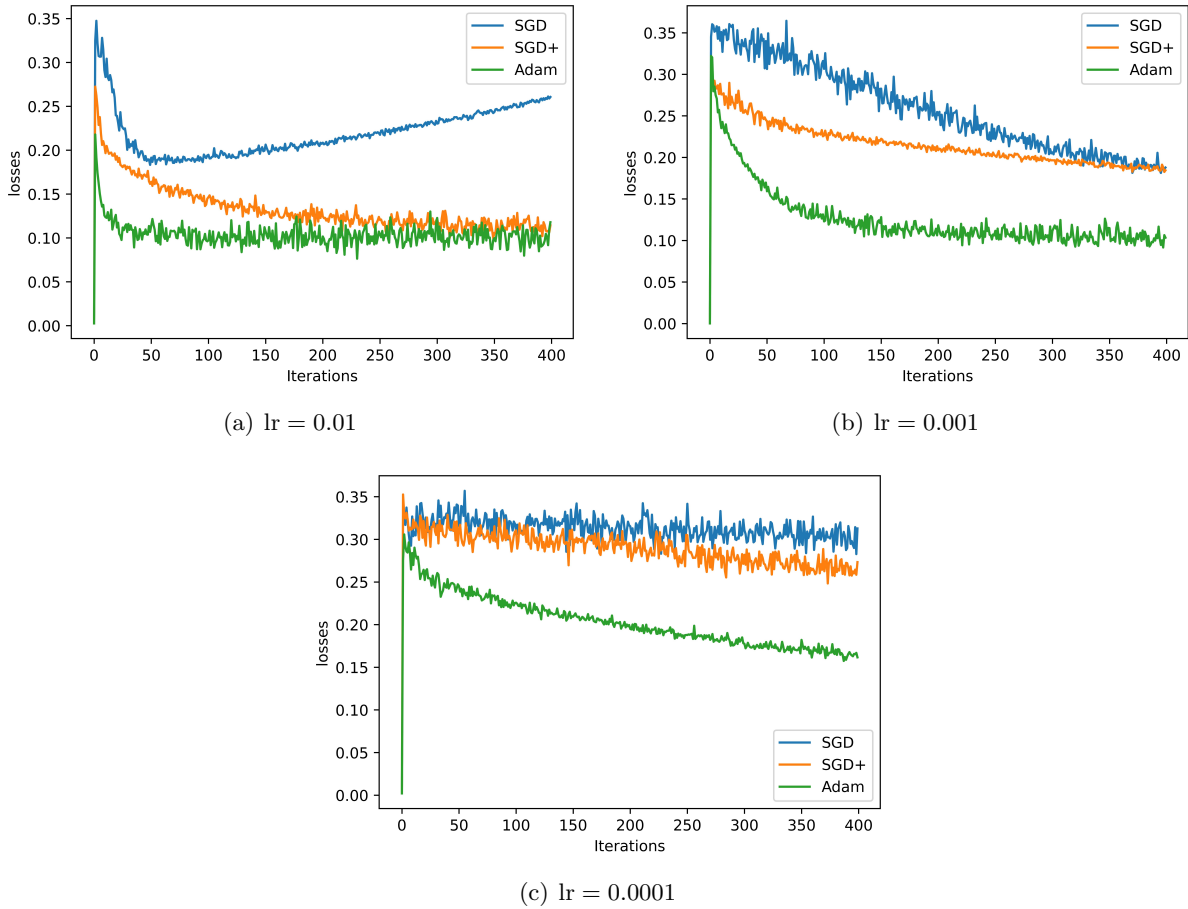


Figure 3: Loss curves for one-neuron case with different learning rates.

2.2 Multi-neuron classifier

Note: The code for this subsection is borrowed from the `ComputationalGraphPrimer` package (<https://engineering.purdue.edu/kak/distCGP/>).

The code for this task is shown below. To make the code concise, I only put the comments in the lines that I modified. The `MultiNeuronSGDCGP` class is for Vanilla Gradient Descent. The `run_training_loop_one_neuron_model` method is rewritten to return the loss record. The `MultiNeuronSGDPlusCGP` class is for SGD+. A new parameter, momentum factor, is added in the `__init__` method. The `run_training_loop_one_neuron_model` method is rewritten to return the loss record. The `backprop_and_update_params_one_neuron_model` method is rewritten to apply the SGD+ weight update equations. The `MultiNeuronAdamCGP` class is for Adam. Three new parameters, first exponential decay rate, second exponential decay rate, and epsilon, are added in the `__init__` method. The `run_training_loop_one_neuron_model` method is rewritten to return the loss record. The `backprop_and_update_params_one_neuron_model` method is rewritten to apply the Adam weight update equations.

```
...
class MultiNeuronSGDCGP(ComputationalGraphPrimer):
    def __init__(self, *args, **kwargs):
        super().__init__(*args, **kwargs) ## inherit from CGP
    def run_training_loop_multi_neuron_model(self, training_data):
        class DataLoader:
            def __init__(self, training_data, batch_size):
                self.training_data = training_data
                self.batch_size = batch_size
                self.class_0_samples = [(item, 0) for item in self.
                    ↪ training_data[0]]
                self.class_1_samples = [(item, 1) for item in self.
                    ↪ training_data[1]]
            def __len__(self):
                return len(self.training_data[0]) + len(self.
                    ↪ training_data[1])
            def _getitem(self):
                cointoss = random.choice([0,1])
                if cointoss == 0:
                    return random.choice(self.class_0_samples)
                else:
                    return random.choice(self.class_1_samples)
            def getbatch(self):
                batch_data, batch_labels = [], []
                maxval = 0.0
                for _ in range(self.batch_size):
                    item = self._getitem()
                    if np.max(item[0]) > maxval:
                        maxval = np.max(item[0])
                    batch_data.append(item[0])
                    batch_labels.append(item[1])
                batch_data = [item/maxval for item in batch_data]
                batch = [batch_data, batch_labels]
                return batch
        self.vals_for_learnable_params = {param: random.uniform(0,1)
            ↪ for param in self.learnable_params}
        self.bias = {i : [random.uniform(0,1) for j in range(self.
            ↪ layers_config[i])] for i in range(1, self.num_layers)}
        data_loader = DataLoader(training_data, batch_size=self.
            ↪ batch_size)
        loss_running_record = []
```

```

i = 0
avg_loss_over_iterations = 0.0
for i in range(self.training_iterations):
    data = data_loader.getbatch()
    data_tuples = data[0]
    class_labels = data[1]
    self.forward_prop_multi_neuron_model(data_tuples)
    predicted_labels_for_batch = self.forw_prop_vals_at_layers[
        ↪ self.num_layers-1]
    y_preds = [item for sublist in predicted_labels_for_batch
        ↪ for item in sublist]
    loss = sum([(abs(class_labels[i] - y_preds[i]))**2 for i in
        ↪ range(len(class_labels))])
    loss_avg = loss / float(len(class_labels))
    avg_loss_over_iterations += loss_avg
    if i%(self.display_loss_how_often) == 0:
        avg_loss_over_iterations /= self.display_loss_how_often
        loss_running_record.append(avg_loss_over_iterations)
        print("[iter=%d] loss = %.4f" % (i+1,
            ↪ avg_loss_over_iterations))
        avg_loss_over_iterations = 0.0
    y_errors_in_batch = list(map(operator.sub, class_labels,
        ↪ y_preds))
    self.backprop_and_update_params_multi_neuron_model(y_preds,
        ↪ y_errors_in_batch)
    ## removed plt functions
    return loss_running_record ## return the loss record

class MultiNeuronSGDPlusCGP(ComputationalGraphPrimer):
    def __init__(self, mu, *args, **kwargs):
        super().__init__(*args, **kwargs) ## inherit from CGP
        self.mu = mu ## momentum factor
    def run_training_loop_multi_neuron_model(self, training_data):
        class DataLoader:
            def __init__(self, training_data, batch_size):
                self.training_data = training_data
                self.batch_size = batch_size
                self.class_0_samples = [(item, 0) for item in self.
                    ↪ training_data[0]]
                self.class_1_samples = [(item, 1) for item in self.
                    ↪ training_data[1]]
            def __len__(self):
                return len(self.training_data[0]) + len(self.
                    ↪ training_data[1])
            def _getitem(self):
                cointoss = random.choice([0,1])
                if cointoss == 0:
                    return random.choice(self.class_0_samples)
                else:
                    return random.choice(self.class_1_samples)
            def getbatch(self):
                batch_data, batch_labels = [], []
                maxval = 0.0
                for _ in range(self.batch_size):

```

```

        item = self._getitem()
        if np.max(item[0]) > maxval:
            maxval = np.max(item[0])
        batch_data.append(item[0])
        batch_labels.append(item[1])
    batch_data = [item/maxval for item in batch_data]
    batch = [batch_data, batch_labels]
    return batch

self.vals_for_learnable_params = {param: random.uniform(0,1)
    ↪ for param in self.learnable_params}
self.bias = {i : [random.uniform(0,1) for j in range(self.
    ↪ layers_config[i])] for i in range(1, self.num_layers)}
self.previous_v_vals = {param: 0 for param in self.
    ↪ learnable_params} ## initialize v_0 for weights
self.previous_v_bias = {i : [0 for j in range(self.
    ↪ layers_config[i])] for i in range(1, self.num_layers)} ##
    ↪ initialize v_0 for bias
data_loader = DataLoader(training_data, batch_size=self.
    ↪ batch_size)
loss_running_record = []
i = 0
avg_loss_over_iterations = 0.0
for i in range(self.training_iterations):
    data = data_loader.getbatch()
    data_tuples = data[0]
    class_labels = data[1]
    self.forward_prop_multi_neuron_model(data_tuples)
    predicted_labels_for_batch = self.forw_prop_vals_at_layers[
        ↪ self.num_layers-1]
    y_preds = [item for sublist in predicted_labels_for_batch
        ↪ for item in sublist]
    loss = sum([(abs(class_labels[i] - y_preds[i]))**2 for i in
        ↪ range(len(class_labels))])
    loss_avg = loss / float(len(class_labels))
    avg_loss_over_iterations += loss_avg
    if i%(self.display_loss_how_often) == 0:
        avg_loss_over_iterations /= self.display_loss_how_often
        loss_running_record.append(avg_loss_over_iterations)
        print("[iter=%d] loss = %.4f" % (i+1,
            ↪ avg_loss_over_iterations))
        avg_loss_over_iterations = 0.0
    y_errors_in_batch = list(map(operator.sub, class_labels,
        ↪ y_preds))
    self.backprop_and_update_params_multi_neuron_model(y_preds,
        ↪ y_errors_in_batch)
## removed plt functions
return loss_running_record ## return the loss record
def backprop_and_update_params_multi_neuron_model(self, predictions
    ↪ , y_errors):
    pred_err_backproped_at_layers = [{i : [None for j in range(self
        ↪ .layers_config[i])] for i in range(self.num_layers)} for
        ↪ _ in range(self.batch_size)]
    partial_of_loss_wrt_params = {param : 0.0 for param in self.
        ↪ all_params}

```

```

bias_changes = {i : [0.0 for j in range(self.layers_config[i
    ↪ ])] for i in range(1, self.num_layers)}
for b in range(self.batch_size):
    pred_err_backproped_at_layers[b][self.num_layers - 1] = [
    ↪ y_errors[b] ]
    for back_layer_index in reversed(range(1, self.num_layers)):
        input_vals = self.forw_prop_vals_at_layers[
            ↪ back_layer_index - 1]
        deriv_sigmoids = self.gradient_vals_for_layers[
            ↪ back_layer_index]
        vars_in_layer = self.layer_vars[back_layer_index]
        vars_in_next_layer_back = self.layer_vars[
            ↪ back_layer_index - 1]
        vals_for_input_vars_dict = dict(zip(
            ↪ vars_in_next_layer_back, self.
            ↪ forw_prop_vals_at_layers[back_layer_index - 1][b
            ↪ ]))
        layer_params = self.layer_params[back_layer_index]
        transposed_layer_params = list(zip(*layer_params))
        for k, var1 in enumerate(vars_in_next_layer_back):
            for j, var2 in enumerate(vars_in_layer):
                pred_err_backproped_at_layers[b][
                    ↪ back_layer_index - 1][k] = sum([self.
                    ↪ vals_for_learnable_params[
                    ↪ transposed_layer_params[k][i]] \
                    * pred_err_backproped_at_layers[b][
                    ↪ back_layer_index][i] for i in range(
                    ↪ len(vars_in_layer))])
        for j, var in enumerate(vars_in_layer):
            layer_params = self.layer_params[back_layer_index][
                ↪ j]
            input_vars_to_param_map = self.var_to_var_param[var
                ↪ ]
            param_to_vars_map = {param : var for var, param in
                ↪ input_vars_to_param_map.items()}
            for i, param in enumerate(layer_params):
                partial_of_loss_wrt_params[param] +=
                    ↪ pred_err_backproped_at_layers[b][
                    ↪ back_layer_index][j] * \
                    vals_for_input_vars_dict[param_to_vars_map[
                        ↪ param]] * deriv_sigmoids[b][j]
            for k, var1 in enumerate(vars_in_next_layer_back):
                for j, var2 in enumerate(vars_in_layer):
                    if back_layer_index - 1 > 0:
                        bias_changes[back_layer_index - 1][k] +=
                            ↪ pred_err_backproped_at_layers[b][
                            ↪ back_layer_index - 1][k] *
                            ↪ deriv_sigmoids[b][j]
        for param in partial_of_loss_wrt_params:
            partial_of_loss_wrt_param = partial_of_loss_wrt_params[
                ↪ param] / float(self.batch_size)
            self.previous_v_vals[param] = self.mu * self.
                ↪ previous_v_vals[param] + partial_of_loss_wrt_param ##
            ↪ update momentum memory

```



```

        step = self.learning_rate * self.previous_v_vals[param] ##
        ↪ scaled step size
        self.vals_for_learnable_params[param] += step
    for layer_index in range(1,self.num_layers):
        for k in range(self.layers_config[layer_index]):
            self.previous_v_bias[layer_index][k] = self.mu * self.
            ↪ previous_v_bias[layer_index][k] + (bias_changes[
            ↪ layer_index][k] / float(self.batch_size)) ##
            ↪ update momentum memory
            self.bias[layer_index][k] += self.learning_rate *
            ↪ self.previous_v_bias[layer_index][k] ## Update
            ↪ the bias

class MultiNeuronAdamCGP(ComputationalGraphPrimer):
    def __init__(self, beta_1, beta_2, epsilon, *args, **kwargs):
        super().__init__(*args, **kwargs) ## inherit from CGP
        self.beta_1 = beta_1 ## first exponential decay rate
        self.beta_2 = beta_2 ## second exponential decay rate
        self.epsilon = epsilon ## epsilon
    def run_training_loop_multi_neuron_model(self, training_data):
        class DataLoader:
            def __init__(self, training_data, batch_size):
                self.training_data = training_data
                self.batch_size = batch_size
                self.class_0_samples = [(item, 0) for item in self.
                ↪ training_data[0]]
                self.class_1_samples = [(item, 1) for item in self.
                ↪ training_data[1]]
            def __len__(self):
                return len(self.training_data[0]) + len(self.
                ↪ training_data[1])
            def _getitem(self):
                cointoss = random.choice([0,1])
                if cointoss == 0:
                    return random.choice(self.class_0_samples)
                else:
                    return random.choice(self.class_1_samples)
            def getbatch(self):
                batch_data, batch_labels = [], []
                maxval = 0.0
                for _ in range(self.batch_size):
                    item = self._getitem()
                    if np.max(item[0]) > maxval:
                        maxval = np.max(item[0])
                    batch_data.append(item[0])
                    batch_labels.append(item[1])
                batch_data = [item/maxval for item in batch_data]
                batch = [batch_data, batch_labels]
                return batch
        self.vals_for_learnable_params = {param: random.uniform(0,1)
        ↪ for param in self.learnable_params}
        self.bias = {i : [random.uniform(0,1) for j in range( self.
        ↪ layers_config[i] ) ] for i in range(1, self.num_layers)}
        self.previous_m_vals = {param: 0 for param in self.

```

```

        ↪ learnable_params} ## initialize m_0 for weights
self.previous_m_bias = {i : [0 for j in range(self.
        ↪ layers_config[i])] for i in range(1, self.num_layers)} ##
        ↪ initialize m_0 for bias
self.previous_v_vals = {param: 0 for param in self.
        ↪ learnable_params} ## initialize m_0 for weights
self.previous_v_bias = {i : [0 for j in range(self.
        ↪ layers_config[i])] for i in range(1, self.num_layers)} ##
        ↪ initialize m_0 for bias
data_loader = DataLoader(training_data, batch_size=self.
        ↪ batch_size)
loss_running_record = []
i = 0
avg_loss_over_iterations = 0.0
for i in range(self.training_iterations):
    data = data_loader.getbatch()
    data_tuples = data[0]
    class_labels = data[1]
    self.forward_prop_multi_neuron_model(data_tuples)
    predicted_labels_for_batch = self.forw_prop_vals_at_layers[
        ↪ self.num_layers-1]
    y_preds = [item for sublist in predicted_labels_for_batch
        ↪ for item in sublist]
    loss = sum([(abs(class_labels[i] - y_preds[i]))**2 for i in
        ↪ range(len(class_labels))])
    loss_avg = loss / float(len(class_labels))
    avg_loss_over_iterations += loss_avg
    if i%(self.display_loss_how_often) == 0:
        avg_loss_over_iterations /= self.display_loss_how_often
        loss_running_record.append(avg_loss_over_iterations)
        print("[iter=%d] loss = %.4f" % (i+1,
            ↪ avg_loss_over_iterations))
        avg_loss_over_iterations = 0.0
    y_errors_in_batch = list(map(operator.sub, class_labels,
        ↪ y_preds))
    self.backprop_and_update_params_multi_neuron_model(y_preds,
        ↪ y_errors_in_batch, i+1)
## removed plt functions
return loss_running_record ## return the loss record
def backprop_and_update_params_multi_neuron_model(self, predictions
    ↪ , y_errors, iter): ## add iteration index as aparameter
    pred_err_backproped_at_layers = [{i : [None for j in range(self
        ↪ .layers_config[i])] for i in range(self.num_layers)} for
        ↪ _ in range(self.batch_size)]
    partial_of_loss_wrt_params = {param : 0.0 for param in self.
        ↪ all_params}
    bias_changes = {i : [0.0 for j in range(self.layers_config[i
        ↪ ])] for i in range(1, self.num_layers)}
    for b in range(self.batch_size):
        pred_err_backproped_at_layers[b][self.num_layers - 1] = [
            ↪ y_errors[b]]
        for back_layer_index in reversed(range(1, self.num_layers)):
            input_vals = self.forw_prop_vals_at_layers[
                ↪ back_layer_index -1]

```

```

deriv_sigmoids = self.gradient_vals_for_layers[
    ↪ back_layer_index]
vars_in_layer = self.layer_vars[back_layer_index]
vars_in_next_layer_back = self.layer_vars[
    ↪ back_layer_index - 1]
vals_for_input_vars_dict = dict(zip(
    ↪ vars_in_next_layer_back, self.
    ↪ forw_prop_vals_at_layers[back_layer_index - 1][b
    ↪ ]))
layer_params = self.layer_params[back_layer_index]
transposed_layer_params = list(zip(*layer_params))
for k,var1 in enumerate(vars_in_next_layer_back):
    for j,var2 in enumerate(vars_in_layer):
        pred_err_backproped_at_layers[b][
            ↪ back_layer_index - 1][k] = sum([self.
            ↪ vals_for_learnable_params[
            ↪ transposed_layer_params[k][i]] \
            * pred_err_backproped_at_layers[b][
            ↪ back_layer_index][i] for i in range(
            ↪ len(vars_in_layer))])
for j,var in enumerate(vars_in_layer):
    layer_params = self.layer_params[back_layer_index][
        ↪ j]
    input_vars_to_param_map = self.var_to_var_param[var
        ↪ ]
    param_to_vars_map = {param : var for var, param in
        ↪ input_vars_to_param_map.items()}
    for i,param in enumerate(layer_params):
        partial_of_loss_wrt_params[param] += -
            ↪ pred_err_backproped_at_layers[b][
            ↪ back_layer_index][j] * \
            vals_for_input_vars_dict[param_to_vars_map[
            ↪ param]] * deriv_sigmoids[b][j]
    for k,var1 in enumerate(vars_in_next_layer_back):
        for j,var2 in enumerate(vars_in_layer):
            if back_layer_index-1 > 0:
                bias_changes[back_layer_index-1][k] += -
                    ↪ pred_err_backproped_at_layers[b][
                    ↪ back_layer_index - 1][k] *
                    ↪ deriv_sigmoids[b][j]
for param in partial_of_loss_wrt_params:
    partial_of_loss_wrt_param = partial_of_loss_wrt_params[
        ↪ param] / float(self.batch_size)
    self.previous_m_vals[param] = self.beta_1 * self.
        ↪ previous_m_vals[param] + (1 - self.beta_1) *
        ↪ partial_of_loss_wrt_param ## update first momentum
        ↪ memory
    self.previous_v_vals[param] = self.beta_2 * self.
        ↪ previous_v_vals[param] + (1 - self.beta_2) *
        ↪ partial_of_loss_wrt_param**2 ## update second
        ↪ momentum memory
    m_t_hat = self.previous_m_vals[param] / (1 - self.beta_1**
        ↪ iter) ## scaled first momentum memory
    v_t_hat = self.previous_v_vals[param] / (1 - self.beta_2**

```

```

        ↪ iter) ## scaled second momentum memory
    step = self.learning_rate * m_t_hat / np.sqrt(v_t_hat +
        ↪ self.epsilon) ## scaled step size
    self.vals_for_learnable_params[param] -= step
## Finally we update the biases at all the nodes that aggregate
    ↪ data:
for layer_index in range(1,self.num_layers):
    for k in range(self.layers_config[layer_index]):
        self.previous_m_bias[layer_index][k] = self.beta_1 *
            ↪ self.previous_m_bias[layer_index][k] + (1 - self.
            ↪ beta_1) * (bias_changes[layer_index][k] / float(
            ↪ self.batch_size)) ## update first momentum memory
        self.previous_v_bias[layer_index][k] = self.beta_2 *
            ↪ self.previous_v_bias[layer_index][k] + (1 - self.
            ↪ beta_2) * (bias_changes[layer_index][k] / float(
            ↪ self.batch_size))*2 ## update second momentum
            ↪ memory
        m_t_hat = self.previous_m_bias[layer_index][k] / (1 -
            ↪ self.beta_1**iter) ## scaled first momentum
            ↪ memory
        v_t_hat = self.previous_v_bias[layer_index][k] / (1 -
            ↪ self.beta_2**iter) ## scaled second momentum
            ↪ memory
        self.bias[layer_index][k] -= self.learning_rate *
            ↪ m_t_hat / np.sqrt(v_t_hat + self.epsilon) ##
            ↪ Update the bias
...
## run multi-neuron case with different optimizer
def run_multi_neuron_diff_optimizer(lr):
    cgp_SGD = MultiNeuronSGDCGP(
        num_layers = 3,
        layers_config = [4,2,1], # num of nodes in each layer
        expressions = ['xw=ap*xp+aq*xq+ar*xr+as*xs',
                       'xz=bp*xp+bq*xq+br*xr+bs*xs',
                       'xo=cp*xw+cq*xz'],
        output_vars = ['xo'],
        dataset_size = 5000,
        learning_rate = lr,
        training_iterations = 20000,
        batch_size = 8,
        display_loss_how_often = 100,
        debug = True,
    )
    cgp_SGD.parse_multi_layer_expressions()
    cgp_SGD.display_multi_neuron_network()

    cgp_SGDPlus = MultiNeuronSGDPlusCGP(
        mu = 0.5,
        num_layers = 3,
        layers_config = [4,2,1], # num of nodes in each layer
        expressions = ['xw=ap*xp+aq*xq+ar*xr+as*xs',
                       'xz=bp*xp+bq*xq+br*xr+bs*xs',
                       'xo=cp*xw+cq*xz'],
        output_vars = ['xo'],

```

```

        dataset_size = 5000,
        learning_rate = lr,
        training_iterations = 20000,
        batch_size = 8,
        display_loss_how_often = 100,
        debug = True,
    )
    cgp_SGDPlus.parse_multi_layer_expressions()
    cgp_SGDPlus.display_multi_neuron_network()

    cgp_Adam = MultiNeuronAdamCGP(
        beta_1 = 0.9,
        beta_2 = 0.99,
        epsilon = 1e-8,
        num_layers = 3,
        layers_config = [4,2,1], # num of nodes in each layer
        expressions = ['xw=ap*xp+aq*xq+ar*xr+as*xs',
                       'xz=bp*xp+bq*xq+br*xr+bs*xs',
                       'xo=cp*xw+cq*xz'],
        output_vars = ['xo'],
        dataset_size = 5000,
        learning_rate = lr,
        training_iterations = 20000,
        batch_size = 8,
        display_loss_how_often = 100,
        debug = True,
    )
    cgp_Adam.parse_multi_layer_expressions()
    cgp_Adam.display_multi_neuron_network()

    training_data = cgp_SGD.gen_training_data()
    loss_SGD = cgp_SGD.run_training_loop_multi_neuron_model(
        ↪ training_data)
    loss_SGDPlus = cgp_SGDPlus.run_training_loop_multi_neuron_model(
        ↪ training_data)
    loss_Adam = cgp_Adam.run_training_loop_multi_neuron_model(
        ↪ training_data)

    plot_learning_curve(loss_SGD, loss_SGDPlus, loss_Adam, '
        ↪ multi_neuron', str(lr))
...
if __name__ == '__main__':
    ...
    if args.task == '2':
        for lr in [1e-2, 5e-3, 1e-3]:
            run_multi_neuron_diff_optimizer(lr)
    ...

```

Fig. 4 shows the loss curves with the learning rates of 0.01, 0.005, and 0.001 respectively. Adam performs the best with no surprise, and SGD still does the worst. When the learning rate is large (as in Fig. 4(a)), Adam converges fastest, and its loss curve is smooth during the converging, but it oscillates after achieving the global minimum point with the same reason mentioned in the last subsection; SGD+ converges slower with a smooth curve and seems hard to find the optimal point; SGD only converges a little. When the learning rate is moderate (as in Fig. 4(b)), Adam still

converges fastest but a little slower with this smaller learning rate; SGD+ and SGD act similarly at the beginning (both oscillate), but SGD+ converges more after the half training. When the learning rate is small (as in Fig. 3(c)), only Adam converges, but it takes a longer time; SGD+ and SGD both do not converge, oscillating obviously all the way. Compared with Fig. 1(b), `torch.nn` performs better than SGD+ and SGD, but it is not as good as Adam. Although `torch.nn` is optimized well, dealing with more complex problems (more layers), Adam wins with better logic.

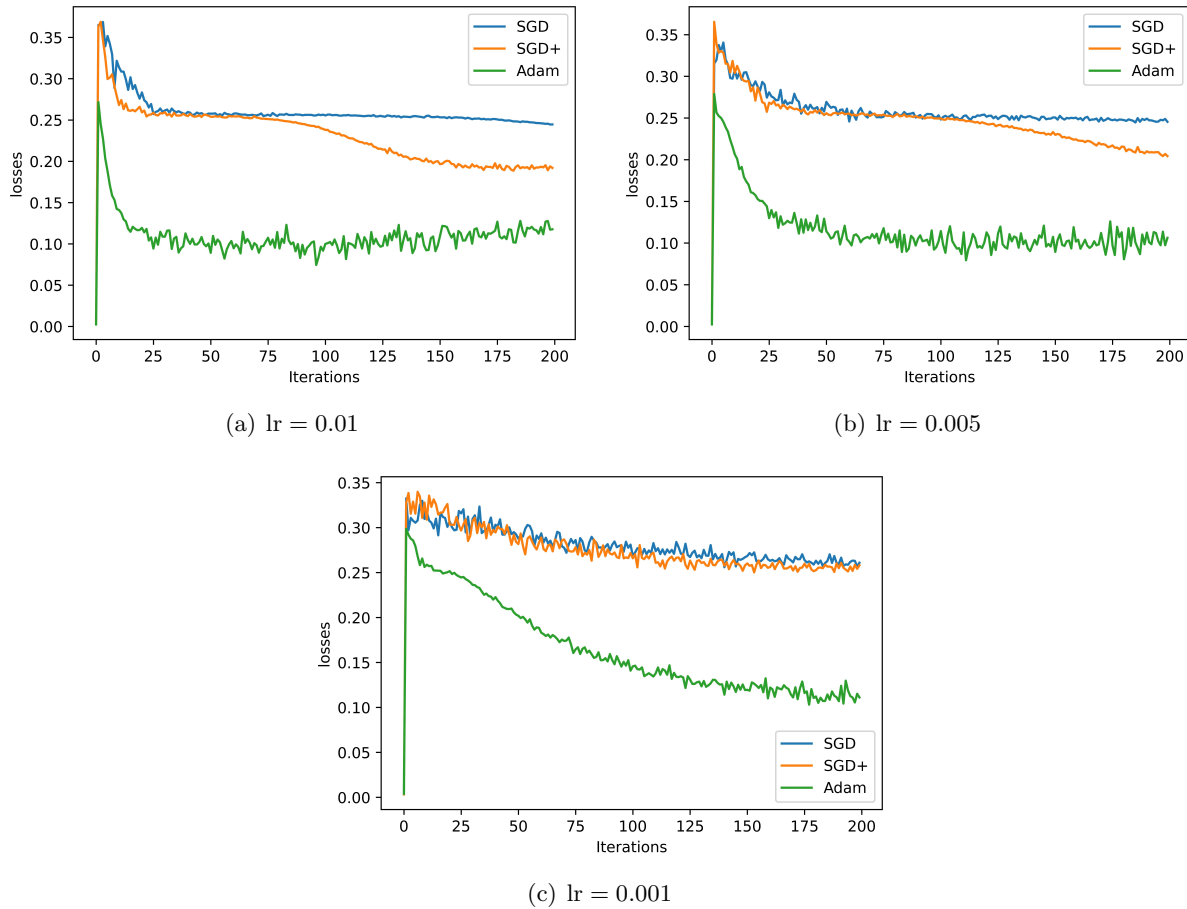


Figure 4: Loss curves for multi-neuron case with different learning rates.

2.3 Exploring hyperparameters of Adam

Note: The code for this subsection is borrowed from the `ComputationalGraphPrimer` package (<https://engineering.purdue.edu/kak/distCGP/>).

The code for this task is shown below. To test different hyperparameters, `run_one_neuron_diff_adam` function is created and called in a `for` loop with different parameters. In this task, only one-neuron model is used.

```
...
## run one-neuron case with adam using different beta_1 and beta_2
def run_one_neuron_diff_adam(beta_1, beta_2):
    cgp_Adam = OneNeuronAdamCGP(
        beta_1 = beta_1,
        beta_2 = beta_2,
```

```

        epsilon = 1e-8,
        one_neuron_model = True,
        expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
        output_vars = ['xw'],
        dataset_size = 5000,
        learning_rate = 1e-3,
        training_iterations = 40000,
        batch_size = 8,
        display_loss_how_often = 100,
        debug = True,
    )
    cgp_Adam.parse_expressions()
    cgp_Adam.display_one_neuron_network()
    training_data = cgp_Adam.gen_training_data()
    start_time = time.time()
    loss_Adam = cgp_Adam.run_training_loop_one_neuron_model(
        ↪ training_data)
    time_used = time.time() - start_time
    final_loss = loss_Adam[-1]
    min_loss = min(loss_Adam[1:])
    return time_used, final_loss, min_loss

if __name__ == '__main__':
    ...
    if args.task == '3':
        beta_list = []
        time_list = []
        final_l_list = []
        min_l_list = []
        for beta_1 in [0.8, 0.9, 0.99]:
            for beta_2 in [0.8, 0.9, 0.99]:
                time_used, final_loss, min_loss =
                    ↪ run_one_neuron_diff_adam(beta_1, beta_2)
                beta_list.append((beta_1, beta_2))
                time_list.append(time_used)
                final_l_list.append(final_loss)
                min_l_list.append(min_loss)
        for i in range(len(beta_list)):
            print('Performance with beta_1 %.2f and beta_2 %.2f \n----
                ↪ time: %fs; final loss: %f; minimum loss: %f'%(
                ↪ beta_list[i][0], beta_list[i][1], time_list[i],
                ↪ final_l_list[i], min_l_list[i]))

```

Output:

```

Performance with beta_1 0.80 and beta_2 0.80
---- time: 8.037810s; final loss: 0.097885; minimum loss: 0.085376
Performance with beta_1 0.80 and beta_2 0.90
---- time: 8.147538s; final loss: 0.107459; minimum loss: 0.084445
Performance with beta_1 0.80 and beta_2 0.99
---- time: 8.146399s; final loss: 0.105753; minimum loss: 0.083214
Performance with beta_1 0.90 and beta_2 0.80
---- time: 8.216441s; final loss: 0.091616; minimum loss: 0.083619
Performance with beta_1 0.90 and beta_2 0.90
---- time: 8.260545s; final loss: 0.090110; minimum loss: 0.079896

```

```

Performance with beta_1 0.90 and beta_2 0.99
---- time: 8.219999s; final loss: 0.085230; minimum loss: 0.079815
Performance with beta_1 0.99 and beta_2 0.80
---- time: 8.225816s; final loss: 0.095684; minimum loss: 0.088200
Performance with beta_1 0.99 and beta_2 0.90
---- time: 8.173840s; final loss: 0.099921; minimum loss: 0.087710
Performance with beta_1 0.99 and beta_2 0.99
---- time: 8.106234s; final loss: 0.112987; minimum loss: 0.081915

```

β_1 is used to maintain the direction of the gradient from the previous step, making the update more stable and helping to get out of local minimum points. β_2 prevents the gradient from being updated too aggressively, which is similar to adding a penalty term. The performance of Adam with different β_1 and β_2 is shown in Tab. 1. The running times are similar, meaning these two parameters do not affect the running time for each iteration. The final losses are various, but it does not matter, since the loss curve oscillates. But the minimum losses are important. When β_1 is a constant, the minimum loss decreases with increasing β_2 . When β_2 is a constant, the minimum loss is smallest with $\beta_1 = 0.9$ (neither too large nor too small!). The best combination is $\beta_1 = 0.9$ and $\beta_2 = 0.99$. These are also the recommended values from the Adam paper <https://arxiv.org/abs/1412.6980>.

Table 1: performance of Adam with different hyperparameters.

β_1	β_2	running time (s)	final loss	minimum loss
0.8	0.8	8.0378	0.0979	0.0854
	0.9	8.1475	0.1075	0.0844
	0.99	8.1464	0.1058	0.0832
0.9	0.8	8.2164	0.0916	0.0836
	0.9	8.2605	0.0901	0.0799
	0.99	8.2200	0.0852	0.0798
0.99	0.8	8.2258	0.0957	0.0882
	0.9	8.1738	0.0999	0.0877
	0.99	8.1062	0.1130	0.0819

3 Source Code

```

1 import sys, os, os.path
2 import numpy as np
3 import re
4 import operator
5 import itertools
6 import math
7 import random
8 import argparse
9 import time
10 import torch
11 from collections import deque
12 import copy
13 import matplotlib.pyplot as plt
14 import networkx as nx
15
16 from ComputationalGraphPrimer import *
17

```



```

18 ## Set random seed
19 def set_random_seed(seed):
20     random.seed(seed)
21     torch.manual_seed(seed)
22     torch.cuda.manual_seed(seed)
23     np.random.seed(seed)
24     torch.backends.cudnn.deterministic = True
25     torch.backends.cudnn.benchmark = False
26     os.environ['PYTHONHASHSEED'] = str(seed)
27
28 class OneNeuronSGDCGP(ComputationalGraphPrimer):
29     def __init__(self, *args, **kwargs):
30         super().__init__(*args, **kwargs) ## inherit from CGP
31     def run_training_loop_one_neuron_model(self, training_data):
32         self.vals_for_learnable_params = {param: random.uniform(0,1)
33             ↪ for param in self.learnable_params}
34         self.bias = random.uniform(0,1)
35         class DataLoader:
36             def __init__(self, training_data, batch_size):
37                 self.training_data = training_data
38                 self.batch_size = batch_size
39                 self.class_0_samples = [(item, 0) for item in self.
40                     ↪ training_data[0]]
41                 self.class_1_samples = [(item, 1) for item in self.
42                     ↪ training_data[1]]
43             def __len__(self):
44                 return len(self.training_data[0]) + len(self.
45                     ↪ training_data[1])
46             def _getitem(self):
47                 cointoss = random.choice([0,1])
48                 if cointoss == 0:
49                     return random.choice(self.class_0_samples)
50                 else:
51                     return random.choice(self.class_1_samples)
52             def getbatch(self):
53                 batch_data, batch_labels = [], []
54                 maxval = 0.0
55                 for _ in range(self.batch_size):
56                     item = self._getitem()
57                     if np.max(item[0]) > maxval:
58                         maxval = np.max(item[0])
59                     batch_data.append(item[0])
60                     batch_labels.append(item[1])
61                 batch_data = [item/maxval for item in batch_data]
62                 batch = [batch_data, batch_labels]
63                 return batch
64         data_loader = DataLoader(training_data, batch_size=self.
65             ↪ batch_size)
66         loss_running_record = []
67         i = 0
68         avg_loss_over_iterations = 0.0
69         for i in range(self.training_iterations):
70             data = data_loader.getbatch()
71             data_tuples_in_batch = data[0]

```

```

67         class_labels_in_batch = data[1]
68         y_preds, deriv_sigmoids = self.
            ↪ forward_prop_one_neuron_model(data_tuples_in_batch)
69         loss = sum([(abs(class_labels_in_batch[i] - y_preds[i]))**2
            ↪ for i in range(len(class_labels_in_batch))])
70         avg_loss_over_iterations += loss / float(len(
            ↪ class_labels_in_batch))
71         if i%(self.display_loss_how_often) == 0:
72             avg_loss_over_iterations /= self.display_loss_how_often
73             loss_running_record.append(avg_loss_over_iterations)
74             print("[iter=%d] loss = %.4f" % (i+1,
            ↪ avg_loss_over_iterations))
75             avg_loss_over_iterations = 0.0
76             y_errors_in_batch = list(map(operator.sub,
            ↪ class_labels_in_batch, y_preds))
77             self.backprop_and_update_params_one_neuron_model(
            ↪ data_tuples_in_batch, y_preds, y_errors_in_batch,
            ↪ deriv_sigmoids)
78         ## removed plt functions
79         return loss_running_record ## return the loss record
80
81 class OneNeuronSGDPlusCGP(ComputationalGraphPrimer):
82     def __init__(self, mu, *args, **kwargs):
83         super().__init__(*args, **kwargs) ## inherit from CGP
84         self.mu = mu ## momentum factor
85     def run_training_loop_one_neuron_model(self, training_data):
86         self.vals_for_learnable_params = {param: random.uniform(0,1)
            ↪ for param in self.learnable_params}
87         self.bias = random.uniform(0,1)
88         self.previous_v_vals = {param: 0 for param in self.
            ↪ learnable_params} ## initialize v_0 for weights
89         self.previous_v_bias = 0 ## initialize v_0 for bias
90         class DataLoader:
91             def __init__(self, training_data, batch_size):
92                 self.training_data = training_data
93                 self.batch_size = batch_size
94                 self.class_0_samples = [(item, 0) for item in self.
            ↪ training_data[0]]
95                 self.class_1_samples = [(item, 1) for item in self.
            ↪ training_data[1]]
96             def __len__(self):
97                 return len(self.training_data[0]) + len(self.
            ↪ training_data[1])
98             def _getitem(self):
99                 cointoss = random.choice([0,1])
100                 if cointoss == 0:
101                     return random.choice(self.class_0_samples)
102                 else:
103                     return random.choice(self.class_1_samples)
104             def getbatch(self):
105                 batch_data, batch_labels = [], []
106                 maxval = 0.0
107                 for _ in range(self.batch_size):
108                     item = self._getitem()

```

```

109         if np.max(item[0]) > maxval:
110             maxval = np.max(item[0])
111             batch_data.append(item[0])
112             batch_labels.append(item[1])
113         batch_data = [item/maxval for item in batch_data]
114         batch = [batch_data, batch_labels]
115         return batch
116     data_loader = DataLoader(training_data, batch_size=self.
117         ↪ batch_size)
118     loss_running_record = []
119     i = 0
120     avg_loss_over_iterations = 0.0
121     for i in range(self.training_iterations):
122         data = data_loader.getbatch()
123         data_tuples_in_batch = data[0]
124         class_labels_in_batch = data[1]
125         y_preds, deriv_sigmoids = self.
126             ↪ forward_prop_one_neuron_model(data_tuples_in_batch)
127         loss = sum([(abs(class_labels_in_batch[i] - y_preds[i]))**2
128             ↪ for i in range(len(class_labels_in_batch))])
129         avg_loss_over_iterations += loss / float(len(
130             ↪ class_labels_in_batch))
131         if i%(self.display_loss_how_often) == 0:
132             avg_loss_over_iterations /= self.display_loss_how_often
133             loss_running_record.append(avg_loss_over_iterations)
134             print("[iter=%d] loss = %.4f" % (i+1,
135                 ↪ avg_loss_over_iterations))
136             avg_loss_over_iterations = 0.0
137             y_errors_in_batch = list(map(operator.sub,
138                 ↪ class_labels_in_batch, y_preds))
139             self.backprop_and_update_params_one_neuron_model(
140                 ↪ data_tuples_in_batch, y_preds, y_errors_in_batch,
141                 ↪ deriv_sigmoids)
142         ## removed plt functions
143         return loss_running_record ## return the loss record
144     def backprop_and_update_params_one_neuron_model(self,
145         ↪ data_tuples_in_batch, predictions, y_errors_in_batch,
146         ↪ deriv_sigmoids):
147         input_vars = self.independent_vars
148         input_vars_to_param_map = self.var_to_var_param[self.
149             ↪ output_vars[0]]
150         param_to_vars_map = {param : var for var, param in
151             ↪ input_vars_to_param_map.items()}
152         vals_for_learnable_params = self.vals_for_learnable_params
153         for i,param in enumerate(self.vals_for_learnable_params):
154             partial_of_loss_wrt_param = 0.0
155             for j in range(self.batch_size):
156                 vals_for_input_vars_dict = dict(zip(input_vars, list(
157                     ↪ data_tuples_in_batch[j])))
158                 partial_of_loss_wrt_param += - y_errors_in_batch[j] *
159                     ↪ vals_for_input_vars_dict[param_to_vars_map[param]
160                     ↪ ] * deriv_sigmoids[j]
161             partial_of_loss_wrt_param /= float(self.batch_size) ##
162                 ↪ average after the batch finishing

```

```

147         self.previous_v_vals[param] = self.mu * self.
            ↳ previous_v_vals[param] + partial_of_loss_wrt_param ##
            ↳ update momentum memory
148         step = self.learning_rate * self.previous_v_vals[param] ##
            ↳ scaled step size
149         self.vals_for_learnable_params[param] -= step
150         partial_of_loss = 0 ## initialize loss count
151         for j in range(self.batch_size): ## loop for a batch
152             partial_of_loss -= y_errors_in_batch[j] * deriv_sigmoids[j]
153         partial_of_loss /= float(self.batch_size) ## average after the
            ↳ batch finishing
154         self.previous_v_bias = self.mu * self.previous_v_bias +
            ↳ partial_of_loss ## update momentum memory
155         self.bias -= self.learning_rate * self.previous_v_bias ##
            ↳ Update the bias
156
157     class OneNeuronAdamCGP(ComputationalGraphPrimer):
158         def __init__(self, beta_1, beta_2, epsilon, *args, **kwargs):
159             super().__init__(*args, **kwargs) ## inherit from CGP
160             self.beta_1 = beta_1 ## first exponential decay rate
161             self.beta_2 = beta_2 ## second exponential decay rate
162             self.epsilon = epsilon ## epsilon
163         def run_training_loop_one_neuron_model(self, training_data):
164             self.vals_for_learnable_params = {param: random.uniform(0,1)
            ↳ for param in self.learnable_params}
165             self.bias = random.uniform(0,1)
166             self.previous_m_vals = {param: 0 for param in self.
            ↳ learnable_params} ## initialize m_0 for weights
167             self.previous_m_bias = 0 ## initialize m_0 for bias
168             self.previous_v_vals = {param: 0 for param in self.
            ↳ learnable_params} ## initialize m_0 for weights
169             self.previous_v_bias = 0 ## initialize m_0 for bias
170         class DataLoader:
171             def __init__(self, training_data, batch_size):
172                 self.training_data = training_data
173                 self.batch_size = batch_size
174                 self.class_0_samples = [(item, 0) for item in self.
            ↳ training_data[0]]
175                 self.class_1_samples = [(item, 1) for item in self.
            ↳ training_data[1]]
176             def __len__(self):
177                 return len(self.training_data[0]) + len(self.
            ↳ training_data[1])
178             def _getitem(self):
179                 cointoss = random.choice([0,1])
180                 if cointoss == 0:
181                     return random.choice(self.class_0_samples)
182                 else:
183                     return random.choice(self.class_1_samples)
184             def getbatch(self):
185                 batch_data, batch_labels = [], []
186                 maxval = 0.0
187                 for _ in range(self.batch_size):
188                     item = self._getitem()

```

```

189         if np.max(item[0]) > maxval:
190             maxval = np.max(item[0])
191             batch_data.append(item[0])
192             batch_labels.append(item[1])
193         batch_data = [item/maxval for item in batch_data]
194         batch = [batch_data, batch_labels]
195         return batch
196     data_loader = DataLoader(training_data, batch_size=self.
197         ↪ batch_size)
198     loss_running_record = []
199     i = 0
200     avg_loss_over_iterations = 0.0
201     for i in range(self.training_iterations):
202         data = data_loader.getbatch()
203         data_tuples_in_batch = data[0]
204         class_labels_in_batch = data[1]
205         y_preds, deriv_sigmoids = self.
206             ↪ forward_prop_one_neuron_model(data_tuples_in_batch)
207         loss = sum([(abs(class_labels_in_batch[i] - y_preds[i]))**2
208             ↪ for i in range(len(class_labels_in_batch))])
209         avg_loss_over_iterations += loss / float(len(
210             ↪ class_labels_in_batch))
211         if i%(self.display_loss_how_often) == 0:
212             avg_loss_over_iterations /= self.display_loss_how_often
213             loss_running_record.append(avg_loss_over_iterations)
214             print("[iter=%d] loss = %.4f" % (i+1,
215                 ↪ avg_loss_over_iterations))
216             avg_loss_over_iterations = 0.0
217             y_errors_in_batch = list(map(operator.sub,
218                 ↪ class_labels_in_batch, y_preds))
219             self.backprop_and_update_params_one_neuron_model(
220                 ↪ data_tuples_in_batch, y_preds, y_errors_in_batch,
221                 ↪ deriv_sigmoids, i+1)
222         ## removed plt functions
223         return loss_running_record ## return the loss record
224     def backprop_and_update_params_one_neuron_model(self,
225         ↪ data_tuples_in_batch, predictions, y_errors_in_batch,
226         ↪ deriv_sigmoids, iter): ## add iteration index as aparameter
227         input_vars = self.independent_vars
228         input_vars_to_param_map = self.var_to_var_param[self.
229             ↪ output_vars[0]]
230         param_to_vars_map = {param : var for var, param in
231             ↪ input_vars_to_param_map.items()}
232         vals_for_learnable_params = self.vals_for_learnable_params
233         for i,param in enumerate(self.vals_for_learnable_params):
234             partial_of_loss_wrt_param = 0.0
235             for j in range(self.batch_size):
236                 vals_for_input_vars_dict = dict(zip(input_vars, list(
237                     ↪ data_tuples_in_batch[j])))
238                 partial_of_loss_wrt_param += - y_errors_in_batch[j] *
239                     ↪ vals_for_input_vars_dict[param_to_vars_map[param]
240                     ↪ ] * deriv_sigmoids[j]
241             partial_of_loss_wrt_param /= float(self.batch_size) ##
242                 ↪ average after the batch finishing

```

```

227         self.previous_m_vals[param] = self.beta_1 * self.
            ↳ previous_m_vals[param] + (1 - self.beta_1) *
            ↳ partial_of_loss_wrt_param ## update first momentum
            ↳ memory
228         self.previous_v_vals[param] = self.beta_2 * self.
            ↳ previous_v_vals[param] + (1 - self.beta_2) *
            ↳ partial_of_loss_wrt_param**2 ## update second
            ↳ momentum memory
229         m_t_hat = self.previous_m_vals[param] / (1 - self.beta_1**
            ↳ iter) ## scaled first momentum memory
230         v_t_hat = self.previous_v_vals[param] / (1 - self.beta_2**
            ↳ iter) ## scaled second momentum memory
231         step = self.learning_rate * m_t_hat / np.sqrt(v_t_hat +
            ↳ self.epsilon) ## scaled step size
232         self.vals_for_learnable_params[param] -= step
233         partial_of_loss = 0 ## initialize loss count
234         for j in range(self.batch_size): ## loop for a batch
235             partial_of_loss -= y_errors_in_batch[j] * deriv_sigmoids[j]
236         partial_of_loss /= float(self.batch_size) ## average after the
            ↳ batch finishing
237         self.previous_m_bias = self.beta_1 * self.previous_m_bias + (1
            ↳ - self.beta_1) * partial_of_loss ## update first momentum
            ↳ memory
238         self.previous_v_bias = self.beta_2 * self.previous_v_bias + (1
            ↳ - self.beta_2) * partial_of_loss**2 ## update second
            ↳ momentum memory
239         m_t_hat = self.previous_m_bias / (1 - self.beta_1**iter) ##
            ↳ scaled first momentum memory
240         v_t_hat = self.previous_v_bias / (1 - self.beta_2**iter) ##
            ↳ scaled second momentum memory
241         self.bias -= self.learning_rate * m_t_hat / np.sqrt(v_t_hat +
            ↳ self.epsilon) ## Update the bias
242
243     class MultiNeuronSGDCGP(ComputationalGraphPrimer):
244         def __init__(self, *args, **kwargs):
245             super().__init__(*args, **kwargs) ## inherit from CGP
246         def run_training_loop_multi_neuron_model(self, training_data):
247             class DataLoader:
248                 def __init__(self, training_data, batch_size):
249                     self.training_data = training_data
250                     self.batch_size = batch_size
251                     self.class_0_samples = [(item, 0) for item in self.
                        ↳ training_data[0]]
252                     self.class_1_samples = [(item, 1) for item in self.
                        ↳ training_data[1]]
253                 def __len__(self):
254                     return len(self.training_data[0]) + len(self.
                        ↳ training_data[1])
255                 def __getitem__(self):
256                     cointoss = random.choice([0,1])
257                     if cointoss == 0:
258                         return random.choice(self.class_0_samples)
259                     else:
260                         return random.choice(self.class_1_samples)

```

```

261         def getbatch(self):
262             batch_data, batch_labels = [], []
263             maxval = 0.0
264             for _ in range(self.batch_size):
265                 item = self._getitem()
266                 if np.max(item[0]) > maxval:
267                     maxval = np.max(item[0])
268                     batch_data.append(item[0])
269                     batch_labels.append(item[1])
270             batch_data = [item/maxval for item in batch_data]
271             batch = [batch_data, batch_labels]
272             return batch
273         self.vals_for_learnable_params = {param: random.uniform(0,1)
274             ↪ for param in self.learnable_params}
275         self.bias = {i : [random.uniform(0,1) for j in range(self.
276             ↪ layers_config[i])] for i in range(1, self.num_layers)}
277         data_loader = DataLoader(training_data, batch_size=self.
278             ↪ batch_size)
279         loss_running_record = []
280         i = 0
281         avg_loss_over_iterations = 0.0
282         for i in range(self.training_iterations):
283             data = data_loader.getbatch()
284             data_tuples = data[0]
285             class_labels = data[1]
286             self.forward_prop_multi_neuron_model(data_tuples)
287             predicted_labels_for_batch = self.forw_prop_vals_at_layers[
288                 ↪ self.num_layers-1]
289             y_preds = [item for sublist in predicted_labels_for_batch
290                 ↪ for item in sublist]
291             loss = sum([(abs(class_labels[i] - y_preds[i]))**2 for i in
292                 ↪ range(len(class_labels))])
293             loss_avg = loss / float(len(class_labels))
294             avg_loss_over_iterations += loss_avg
295             if i%(self.display_loss_how_often) == 0:
296                 avg_loss_over_iterations /= self.display_loss_how_often
297                 loss_running_record.append(avg_loss_over_iterations)
298                 print("[iter=%d] loss = %.4f" % (i+1,
299                     ↪ avg_loss_over_iterations))
300                 avg_loss_over_iterations = 0.0
301             y_errors_in_batch = list(map(operator.sub, class_labels,
302                 ↪ y_preds))
303             self.backprop_and_update_params_multi_neuron_model(y_preds,
304                 ↪ y_errors_in_batch)
305             ## removed plt functions
306             return loss_running_record ## return the loss record
307
308 class MultiNeuronSGDPlusCGP(ComputationalGraphPrimer):
309     def __init__(self, mu, *args, **kwargs):
310         super().__init__(*args, **kwargs) ## inherit from CGP
311         self.mu = mu ## momentum factor
312     def run_training_loop_multi_neuron_model(self, training_data):
313         class DataLoader:
314             def __init__(self, training_data, batch_size):

```



```

306         self.training_data = training_data
307         self.batch_size = batch_size
308         self.class_0_samples = [(item, 0) for item in self.
    ↪ training_data[0]]
309         self.class_1_samples = [(item, 1) for item in self.
    ↪ training_data[1]]
310     def __len__(self):
311         return len(self.training_data[0]) + len(self.
    ↪ training_data[1])
312     def _getitem(self):
313         cointoss = random.choice([0,1])
314         if cointoss == 0:
315             return random.choice(self.class_0_samples)
316         else:
317             return random.choice(self.class_1_samples)
318     def getbatch(self):
319         batch_data, batch_labels = [], []
320         maxval = 0.0
321         for _ in range(self.batch_size):
322             item = self._getitem()
323             if np.max(item[0]) > maxval:
324                 maxval = np.max(item[0])
325             batch_data.append(item[0])
326             batch_labels.append(item[1])
327         batch_data = [item/maxval for item in batch_data]
328         batch = [batch_data, batch_labels]
329         return batch
330     self.vals_for_learnable_params = {param: random.uniform(0,1)
    ↪ for param in self.learnable_params}
331     self.bias = {i : [random.uniform(0,1) for j in range(self.
    ↪ layers_config[i])] for i in range(1, self.num_layers)}
332     self.previous_v_vals = {param: 0 for param in self.
    ↪ learnable_params} ## initialize v_0 for weights
333     self.previous_v_bias = {i : [0 for j in range(self.
    ↪ layers_config[i])] for i in range(1, self.num_layers)} ##
    ↪ initialize v_0 for bias
334     data_loader = DataLoader(training_data, batch_size=self.
    ↪ batch_size)
335     loss_running_record = []
336     i = 0
337     avg_loss_over_iterations = 0.0
338     for i in range(self.training_iterations):
339         data = data_loader.getbatch()
340         data_tuples = data[0]
341         class_labels = data[1]
342         self.forward_prop_multi_neuron_model(data_tuples)
343         predicted_labels_for_batch = self.forw_prop_vals_at_layers[
    ↪ self.num_layers-1]
344         y_preds = [item for sublist in predicted_labels_for_batch
    ↪ for item in sublist]
345         loss = sum([(abs(class_labels[i] - y_preds[i]))**2 for i in
    ↪ range(len(class_labels))])
346         loss_avg = loss / float(len(class_labels))
347         avg_loss_over_iterations += loss_avg

```



```

348         if i%(self.display_loss_how_often) == 0:
349             avg_loss_over_iterations /= self.display_loss_how_often
350             loss_running_record.append(avg_loss_over_iterations)
351             print("[iter=%d] loss = %.4f" % (i+1,
352                 ↪ avg_loss_over_iterations))
353             avg_loss_over_iterations = 0.0
354             y_errors_in_batch = list(map(operator.sub, class_labels,
355                 ↪ y_preds))
356             self.backprop_and_update_params_multi_neuron_model(y_preds,
357                 ↪ y_errors_in_batch)
358             ## removed plt functions
359             return loss_running_record ## return the loss record
360 def backprop_and_update_params_multi_neuron_model(self, predictions
361     ↪ , y_errors):
362     pred_err_backproped_at_layers = [{i : [None for j in range(self
363         ↪ .layers_config[i])]} for i in range(self.num_layers)] for
364         ↪ _ in range(self.batch_size)]
365     partial_of_loss_wrt_params = {param : 0.0 for param in self.
366         ↪ all_params}
367     bias_changes = {i : [0.0 for j in range(self.layers_config[i
368         ↪ ])] for i in range(1, self.num_layers)}
369     for b in range(self.batch_size):
370         pred_err_backproped_at_layers[b][self.num_layers - 1] = [
371             ↪ y_errors[b] ]
372         for back_layer_index in reversed(range(1, self.num_layers)):
373             input_vals = self.forw_prop_vals_at_layers[
374                 ↪ back_layer_index - 1]
375             deriv_sigmoids = self.gradient_vals_for_layers[
376                 ↪ back_layer_index]
377             vars_in_layer = self.layer_vars[back_layer_index]
378             vars_in_next_layer_back = self.layer_vars[
379                 ↪ back_layer_index - 1]
380             vals_for_input_vars_dict = dict(zip(
381                 ↪ vars_in_next_layer_back, self.
382                 ↪ forw_prop_vals_at_layers[back_layer_index - 1][b
383                 ↪ ])
384             layer_params = self.layer_params[back_layer_index]
385             transposed_layer_params = list(zip(*layer_params))
386             for k, var1 in enumerate(vars_in_next_layer_back):
387                 for j, var2 in enumerate(vars_in_layer):
388                     pred_err_backproped_at_layers[b][
389                         ↪ back_layer_index - 1][k] = sum([self.
390                         ↪ vals_for_learnable_params[
391                         ↪ transposed_layer_params[k][i]] \
392                         ↪ * pred_err_backproped_at_layers[b][
393                         ↪ back_layer_index][i] for i in range(
394                         ↪ len(vars_in_layer))])
395             for j, var in enumerate(vars_in_layer):
396                 layer_params = self.layer_params[back_layer_index][
397                     ↪ j]
398                 input_vars_to_param_map = self.var_to_var_param[var
399                     ↪ ]
400                 param_to_vars_map = {param : var for var, param in
401                     ↪ input_vars_to_param_map.items()}

```

```

379         for i,param in enumerate(layer_params):
380             partial_of_loss_wrt_params[param] +=
                 ↳ pred_err_backproped_at_layers[b][
                 ↳ back_layer_index][j] * \
381                 vals_for_input_vars_dict[param_to_vars_map[
                 ↳ param]] * deriv_sigmoids[b][j]
382         for k,var1 in enumerate(vars_in_next_layer_back):
383             for j,var2 in enumerate(vars_in_layer):
384                 if back_layer_index-1 > 0:
385                     bias_changes[back_layer_index-1][k] +=
                         ↳ pred_err_backproped_at_layers[b][
                         ↳ back_layer_index - 1][k] *
                         ↳ deriv_sigmoids[b][j]
386     for param in partial_of_loss_wrt_params:
387         partial_of_loss_wrt_param = partial_of_loss_wrt_params[
                 ↳ param] / float(self.batch_size)
388         self.previous_v_vals[param] = self.mu * self.
                 ↳ previous_v_vals[param] + partial_of_loss_wrt_param ##
                 ↳ update momentum memory
389         step = self.learning_rate * self.previous_v_vals[param] ##
                 ↳ scaled step size
390         self.vals_for_learnable_params[param] += step
391     for layer_index in range(1,self.num_layers):
392         for k in range(self.layers_config[layer_index]):
393             self.previous_v_bias[layer_index][k] = self.mu * self.
                 ↳ previous_v_bias[layer_index][k] + (bias_changes[
                 ↳ layer_index][k] / float(self.batch_size)) ##
                 ↳ update momentum memory
394             self.bias[layer_index][k] += self.learning_rate *
                 ↳ self.previous_v_bias[layer_index][k] ## Update
                 ↳ the bias
395
396 class MultiNeuronAdamCGP(ComputationalGraphPrimer):
397     def __init__(self, beta_1, beta_2, epsilon, *args, **kwargs):
398         super().__init__(*args, **kwargs) ## inherit from CGP
399         self.beta_1 = beta_1 ## first exponential decay rate
400         self.beta_2 = beta_2 ## second exponential decay rate
401         self.epsilon = epsilon ## epsilon
402     def run_training_loop_multi_neuron_model(self, training_data):
403         class DataLoader:
404             def __init__(self, training_data, batch_size):
405                 self.training_data = training_data
406                 self.batch_size = batch_size
407                 self.class_0_samples = [(item, 0) for item in self.
                 ↳ training_data[0]]
408                 self.class_1_samples = [(item, 1) for item in self.
                 ↳ training_data[1]]
409             def __len__(self):
410                 return len(self.training_data[0]) + len(self.
                 ↳ training_data[1])
411             def _getitem(self):
412                 cointoss = random.choice([0,1])
413                 if cointoss == 0:
414                     return random.choice(self.class_0_samples)

```

```

415         else:
416             return random.choice(self.class_1_samples)
417     def getbatch(self):
418         batch_data, batch_labels = [], []
419         maxval = 0.0
420         for _ in range(self.batch_size):
421             item = self._getitem()
422             if np.max(item[0]) > maxval:
423                 maxval = np.max(item[0])
424                 batch_data.append(item[0])
425                 batch_labels.append(item[1])
426         batch_data = [item/maxval for item in batch_data]
427         batch = [batch_data, batch_labels]
428         return batch
429     self.vals_for_learnable_params = {param: random.uniform(0,1)
430         ↪ for param in self.learnable_params}
431     self.bias = {i : [random.uniform(0,1) for j in range( self.
432         ↪ layers_config[i] ) ] for i in range(1, self.num_layers)}
433     self.previous_m_vals = {param: 0 for param in self.
434         ↪ learnable_params} ## initialize m_0 for weights
435     self.previous_m_bias = {i : [0 for j in range(self.
436         ↪ layers_config[i])] for i in range(1, self.num_layers)} ##
437         ↪ initialize m_0 for bias
438     self.previous_v_vals = {param: 0 for param in self.
439         ↪ learnable_params} ## initialize m_0 for weights
440     self.previous_v_bias = {i : [0 for j in range(self.
441         ↪ layers_config[i])] for i in range(1, self.num_layers)} ##
442         ↪ initialize m_0 for bias
443     data_loader = DataLoader(training_data, batch_size=self.
444         ↪ batch_size)
445     loss_running_record = []
446     i = 0
447     avg_loss_over_iterations = 0.0
448     for i in range(self.training_iterations):
449         data = data_loader.getbatch()
450         data_tuples = data[0]
451         class_labels = data[1]
452         self.forward_prop_multi_neuron_model(data_tuples)
453         predicted_labels_for_batch = self.forw_prop_vals_at_layers[
454             ↪ self.num_layers-1]
455         y_preds = [item for sublist in predicted_labels_for_batch
456             ↪ for item in sublist]
457         loss = sum([(abs(class_labels[i] - y_preds[i]))**2 for i in
458             ↪ range(len(class_labels))])
459         loss_avg = loss / float(len(class_labels))
460         avg_loss_over_iterations += loss_avg
461         if i%(self.display_loss_how_often) == 0:
462             avg_loss_over_iterations /= self.display_loss_how_often
463             loss_running_record.append(avg_loss_over_iterations)
464             print("[iter=%d] loss = %.4f" % (i+1,
465                 ↪ avg_loss_over_iterations))
466             avg_loss_over_iterations = 0.0
467         y_errors_in_batch = list(map(operator.sub, class_labels,
468             ↪ y_preds))

```

```

455         self.backprop_and_update_params_multi_neuron_model(y_preds,
456                     ↪ y_errors_in_batch, i+1)
457     ## removed plt functions
458     return loss_running_record ## return the loss record
459 def backprop_and_update_params_multi_neuron_model(self, predictions
460     ↪ , y_errors, iter): ## add iteration index as aparameter
461     pred_err_backproped_at_layers = [{i : [None for j in range(self
462         ↪ .layers_config[i])]} for i in range(self.num_layers)] for
463     ↪ _ in range(self.batch_size)]
464     partial_of_loss_wrt_params = {param : 0.0 for param in self.
465         ↪ all_params}
466     bias_changes = {i : [0.0 for j in range(self.layers_config[i
467         ↪ ])] for i in range(1, self.num_layers)}
468     for b in range(self.batch_size):
469         pred_err_backproped_at_layers[b][self.num_layers - 1] = [
470             ↪ y_errors[b]]
471         for back_layer_index in reversed(range(1, self.num_layers)):
472             input_vals = self.forw_prop_vals_at_layers[
473                 ↪ back_layer_index - 1]
474             deriv_sigmoids = self.gradient_vals_for_layers[
475                 ↪ back_layer_index]
476             vars_in_layer = self.layer_vars[back_layer_index]
477             vars_in_next_layer_back = self.layer_vars[
478                 ↪ back_layer_index - 1]
479             vals_for_input_vars_dict = dict(zip(
480                 ↪ vars_in_next_layer_back, self.
481                 ↪ forw_prop_vals_at_layers[back_layer_index - 1][b
482                 ↪ ])
483             layer_params = self.layer_params[back_layer_index]
484             transposed_layer_params = list(zip(*layer_params))
485             for k, var1 in enumerate(vars_in_next_layer_back):
486                 for j, var2 in enumerate(vars_in_layer):
487                     pred_err_backproped_at_layers[b][
488                         ↪ back_layer_index - 1][k] = sum([self.
489                         ↪ vals_for_learnable_params[
490                         ↪ transposed_layer_params[k][i]] \
491                         ↪ * pred_err_backproped_at_layers[b][
492                         ↪ back_layer_index][i] for i in range(
493                         ↪ len(vars_in_layer))])
494             for j, var in enumerate(vars_in_layer):
495                 layer_params = self.layer_params[back_layer_index][
496                     ↪ j]
497                 input_vars_to_param_map = self.var_to_var_param[var
498                     ↪ ]
499                 param_to_vars_map = {param : var for var, param in
500                     ↪ input_vars_to_param_map.items()}
501                 for i, param in enumerate(layer_params):
502                     partial_of_loss_wrt_params[param] += -
503                         ↪ pred_err_backproped_at_layers[b][
504                         ↪ back_layer_index][j] * \
505                         ↪ vals_for_input_vars_dict[param_to_vars_map[
506                         ↪ param]] * deriv_sigmoids[b][j]
507             for k, var1 in enumerate(vars_in_next_layer_back):
508                 for j, var2 in enumerate(vars_in_layer):

```

```

485         if back_layer_index-1 > 0:
486             bias_changes[back_layer_index-1][k] += -
487                 ↪ pred_err_backproped_at_layers[b][
488                     ↪ back_layer_index - 1][k] *
489                     ↪ deriv_sigmoids[b][j]
490
491     for param in partial_of_loss_wrt_params:
492         partial_of_loss_wrt_param = partial_of_loss_wrt_params[
493             ↪ param] / float(self.batch_size)
494         self.previous_m_vals[param] = self.beta_1 * self.
495             ↪ previous_m_vals[param] + (1 - self.beta_1) *
496             ↪ partial_of_loss_wrt_param ## update first momentum
497             ↪ memory
498         self.previous_v_vals[param] = self.beta_2 * self.
499             ↪ previous_v_vals[param] + (1 - self.beta_2) *
500             ↪ partial_of_loss_wrt_param**2 ## update second
501             ↪ momentum memory
502         m_t_hat = self.previous_m_vals[param] / (1 - self.beta_1**
503             ↪ iter) ## scaled first momentum memory
504         v_t_hat = self.previous_v_vals[param] / (1 - self.beta_2**
505             ↪ iter) ## scaled second momentum memory
506         step = self.learning_rate * m_t_hat / np.sqrt(v_t_hat +
507             ↪ self.epsilon) ## scaled step size
508         self.vals_for_learnable_params[param] -= step
509     ## Finally we update the biases at all the nodes that aggregate
510     ↪ data:
511     for layer_index in range(1, self.num_layers):
512         for k in range(self.layers_config[layer_index]):
513             self.previous_m_bias[layer_index][k] = self.beta_1 *
514                 ↪ self.previous_m_bias[layer_index][k] + (1 - self.
515                 ↪ beta_1) * (bias_changes[layer_index][k] / float(
516                 ↪ self.batch_size)) ## update first momentum memory
517             self.previous_v_bias[layer_index][k] = self.beta_2 *
518                 ↪ self.previous_v_bias[layer_index][k] + (1 - self.
519                 ↪ beta_2) * (bias_changes[layer_index][k] / float(
520                 ↪ self.batch_size))**2 ## update second momentum
521                 ↪ memory
522             m_t_hat = self.previous_m_bias[layer_index][k] / (1 -
523                 ↪ self.beta_1**iter) ## scaled first momentum
524                 ↪ memory
525             v_t_hat = self.previous_v_bias[layer_index][k] / (1 -
526                 ↪ self.beta_2**iter) ## scaled second momentum
527                 ↪ memory
528             self.bias[layer_index][k] -= self.learning_rate *
529                 ↪ m_t_hat / np.sqrt(v_t_hat + self.epsilon) ##
530                 ↪ Update the bias
531
532     ## plot the learning curves
533     def plot_learning_curve(loss_SGD, loss_SGDPlus, loss_Adam, title, lr):
534         plt.figure(figsize=(6, 4))
535         plt.plot(loss_SGD, label='SGD')
536         plt.plot(loss_SGDPlus, label='SGD+')
537         plt.plot(loss_Adam, label='Adam')
538         plt.xlabel('Iterations')
539         plt.ylabel('losses')

```

```

512 plt.legend()
513 cf = plt.gcf()
514 cf.savefig('%s_%s.jpg'%(title,lr), bbox_inches='tight', \
515           format='jpg', dpi=1000)
516
517 ## run one-neuron case with different optimizer
518 def run_one_neuron_diff_optimizer(lr):
519     cgp_SGD = OneNeuronSGDCGP(
520         one_neuron_model = True,
521         expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
522         output_vars = ['xw'],
523         dataset_size = 5000,
524         learning_rate = lr,
525         training_iterations = 40000,
526         batch_size = 8,
527         display_loss_how_often = 100,
528         debug = True,
529     )
530     cgp_SGD.parse_expressions()
531     cgp_SGD.display_one_neuron_network()
532
533     cgp_SGDPlus = OneNeuronSGDPlusCGP(
534         mu = 0.5,
535         one_neuron_model = True,
536         expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
537         output_vars = ['xw'],
538         dataset_size = 5000,
539         learning_rate = lr,
540         training_iterations = 40000,
541         batch_size = 8,
542         display_loss_how_often = 100,
543         debug = True,
544     )
545     cgp_SGDPlus.parse_expressions()
546     cgp_SGDPlus.display_one_neuron_network()
547
548     cgp_Adam = OneNeuronAdamCGP(
549         beta_1 = 0.9,
550         beta_2 = 0.99,
551         epsilon = 1e-8,
552         one_neuron_model = True,
553         expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
554         output_vars = ['xw'],
555         dataset_size = 5000,
556         learning_rate = lr,
557         training_iterations = 40000,
558         batch_size = 8,
559         display_loss_how_often = 100,
560         debug = True,
561     )
562     cgp_Adam.parse_expressions()
563     cgp_Adam.display_one_neuron_network()
564
565     training_data = cgp_SGD.gen_training_data()

```

```

566     loss_SGD = cgp_SGD.run_training_loop_one_neuron_model(training_data
567         ↪ )
568     loss_SGDPlus = cgp_SGDPlus.run_training_loop_one_neuron_model(
569         ↪ training_data)
570     loss_Adam = cgp_Adam.run_training_loop_one_neuron_model(
571         ↪ training_data)
572
573     plot_learning_curve(loss_SGD, loss_SGDPlus, loss_Adam, 'one_neuron'
574         ↪ , str(lr))
575
576 ## run multi-neuron case with different optimizer
577 def run_multi_neuron_diff_optimizer(lr):
578     cgp_SGD = MultiNeuronSGDCGP(
579         num_layers = 3,
580         layers_config = [4,2,1], # num of nodes in each layer
581         expressions = ['xw=ap*xp+aq*xq+ar*xr+as*xs',
582                       'xz=bp*xp+bq*xq+br*xr+bs*xs',
583                       'xo=cp*xw+cq*xz'],
584         output_vars = ['xo'],
585         dataset_size = 5000,
586         learning_rate = lr,
587         training_iterations = 20000,
588         batch_size = 8,
589         display_loss_how_often = 100,
590         debug = True,
591     )
592     cgp_SGD.parse_multi_layer_expressions()
593     cgp_SGD.display_multi_neuron_network()
594
595     cgp_SGDPlus = MultiNeuronSGDPlusCGP(
596         mu = 0.5,
597         num_layers = 3,
598         layers_config = [4,2,1], # num of nodes in each layer
599         expressions = ['xw=ap*xp+aq*xq+ar*xr+as*xs',
600                       'xz=bp*xp+bq*xq+br*xr+bs*xs',
601                       'xo=cp*xw+cq*xz'],
602         output_vars = ['xo'],
603         dataset_size = 5000,
604         learning_rate = lr,
605         training_iterations = 20000,
606         batch_size = 8,
607         display_loss_how_often = 100,
608         debug = True,
609     )
610     cgp_SGDPlus.parse_multi_layer_expressions()
611     cgp_SGDPlus.display_multi_neuron_network()
612
613     cgp_Adam = MultiNeuronAdamCGP(
614         beta_1 = 0.9,
615         beta_2 = 0.99,
616         epsilon = 1e-8,
617         num_layers = 3,
618         layers_config = [4,2,1], # num of nodes in each layer
619         expressions = ['xw=ap*xp+aq*xq+ar*xr+as*xs',

```

```

616         'xz=bp*xp+bq*xq+br*xr+bs*xs',
617         'xo=cp*xw+cq*xz'],
618         output_vars = ['xo'],
619         dataset_size = 5000,
620         learning_rate = lr,
621         training_iterations = 20000,
622         batch_size = 8,
623         display_loss_how_often = 100,
624         debug = True,
625     )
626     cgp_Adam.parse_multi_layer_expressions()
627     cgp_Adam.display_multi_neuron_network()
628
629     training_data = cgp_SGD.gen_training_data()
630     loss_SGD = cgp_SGD.run_training_loop_multi_neuron_model(
631         ↪ training_data)
632     loss_SGDPlus = cgp_SGDPlus.run_training_loop_multi_neuron_model(
633         ↪ training_data)
634     loss_Adam = cgp_Adam.run_training_loop_multi_neuron_model(
635         ↪ training_data)
636
637     plot_learning_curve(loss_SGD, loss_SGDPlus, loss_Adam, '
638         ↪ multi_neuron', str(lr))
639
640     ## run one-neuron case with adam using different beta_1 and beta_2
641     def run_one_neuron_diff_adam(beta_1, beta_2):
642         cgp_Adam = OneNeuronAdamCGP(
643             beta_1 = beta_1,
644             beta_2 = beta_2,
645             epsilon = 1e-8,
646             one_neuron_model = True,
647             expressions = ['xw=ab*xa+bc*xb+cd*xc+ac*xd'],
648             output_vars = ['xw'],
649             dataset_size = 5000,
650             learning_rate = 1e-3,
651             training_iterations = 40000,
652             batch_size = 8,
653             display_loss_how_often = 100,
654             debug = True,
655         )
656         cgp_Adam.parse_expressions()
657         cgp_Adam.display_one_neuron_network()
658         training_data = cgp_Adam.gen_training_data()
659         start_time = time.time()
660         loss_Adam = cgp_Adam.run_training_loop_one_neuron_model(
661             ↪ training_data)
662         time_used = time.time() - start_time
663         final_loss = loss_Adam[-1]
664         min_loss = min(loss_Adam[1:])
665         return time_used, final_loss, min_loss
666
667     if __name__ == '__main__':
668         # '1' -- one-neuron classifier with different learning rate
669         # '2' -- multi-neuron classifier with different learning rate

```



```

665 # '3' -- compare different beta_1 and beta_2 in adam with one-
        ↳ neuron case
666 parser = argparse.ArgumentParser()
667 parser.add_argument('-t', '--task', type=str, default='1',\
668     help='choose a task', choices=['1','2','3'])
669 args = parser.parse_args()
670
671 seed = 0
672 # set_random_seed(seed)
673
674 if args.task == '1':
675     for lr in [1e-2, 1e-3, 1e-4]:
676         run_one_neuron_diff_optimizer(lr)
677
678 if args.task == '2':
679     for lr in [1e-2, 5e-3, 1e-3]:
680         run_multi_neuron_diff_optimizer(lr)
681
682 if args.task == '3':
683     beta_list = []
684     time_list = []
685     final_l_list = []
686     min_l_list = []
687     for beta_1 in [0.8, 0.9, 0.99]:
688         for beta_2 in [0.8, 0.9, 0.99]:
689             time_used, final_loss, min_loss =
                ↳ run_one_neuron_diff_adam(beta_1, beta_2)
690             beta_list.append((beta_1, beta_2))
691             time_list.append(time_used)
692             final_l_list.append(final_loss)
693             min_l_list.append(min_loss)
694     for i in range(len(beta_list)):
695         print('Performance with beta_1 %.2f and beta_2 %.2f \n----
                ↳ time: %fs; final loss: %f; minimum loss: %f'%(
                ↳ beta_list[i][0], beta_list[i][1], time_list[i],
                ↳ final_l_list[i], min_l_list[i]))

```