

ADELTA: Transpilation Between Deep Learning Frameworks

Linyuan Gong, Jiayi Wang, Alvin Cheung

University of California, Berkeley

Abstract

We propose the **Adversarial DEep Learning Transpiler (ADELT)**, a novel approach to source-to-source transpilation between deep learning frameworks. ADELTA uniquely decouples code skeleton transpilation and API keyword mapping. For code skeleton transpilation, it uses few-shot prompting on large language models (LLMs), while for API keyword mapping, it uses contextual embeddings from a code-specific BERT. These embeddings are trained in a domain-adversarial setup to generate a keyword translation dictionary. ADELTA is trained on an unlabeled web-crawled deep learning corpus, without relying on any hand-crafted rules or parallel data. It outperforms state-of-the-art transpilers, improving pass@1 rate by 17.4 pts and 15.0 pts for PyTorch-Keras and PyTorch-MXNet transpilation pairs respectively. We provide open access to our code at <https://github.com/gonglinyuan/adelt>.

1 Introduction

The rapid development of deep learning (DL) has led to an equally fast emergence of new software frameworks for training neural networks. Unfortunately, maintaining a deep learning framework and keeping it up-to-date is not an easy task. Many deep learning frameworks are deprecated or lose popularity every year, and porting deep learning code from a legacy framework to a new one is a tedious and error-prone task. A *source-to-source transpiler between DL frameworks* would greatly help practitioners overcome this difficulty.

Two promising solutions to source-to-source transpilation between deep learning frameworks are unsupervised neural machine translation (NMT) [Artetxe *et al.*, 2018] and large language models (LLMs). NMT treats deep learning code as a sentence for training sequence-to-sequence [Sutskever *et al.*, 2014] models, but its applicability is limited due to the scarcity of parallel corpora and its notable data hunger. On the other hand, LLMs like GPT-3 [Brown *et al.*, 2020], pretrained on web crawl data, offer potential, performing translation tasks in a few-shot or zero-shot manner. Our early experiments with GPT-4 show its potential in few-shot transpilation of

deep learning programs. However, such models struggle with API-specific details, inaccurately handling function names and parameter mappings.

That said, most deep learning framework code is *structured*: each type of layers has its own constructor, and constructing a network involves calling each layer’s constructor in a chaining manner. By leveraging the structures of programming languages, we can *decouple* the transpilation of skeletal codes from the mapping of API keywords. The transpilation of skeletal codes is the easier part, and LLMs already do a great job. We only need a separate algorithm to translate the *API keywords*, i.e., the function and parameter names to complete the transpilation.

In this paper, we present ADELTA (Figure 1), a method motivated by this insight to transpile DL code. The canonicalized source code is decoupled into two parts: the code skeleton and the API keywords. ADELTA transpiles the code skeleton using a pretrained LLM by few-shot prompting. Each API keyword occurrence is then embedded into a vector by PyBERT, a BERT pretrained on Python code. This vector is both the textual and the contextual representation of the API keyword. ADELTA then leverages domain-adversarial training to learn a generator that maps the vector to an aligned embedding space. The alignment is enforced by a two-player game, where a discriminator is trained to distinguish between the embeddings from the source DL framework and those from the target DL framework. The API keyword embeddings are trained jointly with the generator as the output embedding matrix of a softmax classifier on the aligned embedding space. After generating a synthetic API keyword dictionary from the embeddings using a two-step greedy algorithm, ADELTA then looks up each API keyword occurrence in the dictionary and puts them back into the transpiled code skeleton.

In summary, this paper makes the following contributions:

- We introduce ADELTA, a robust solution for transpilation between deep learning frameworks without training on any labeled data. Outperforming large language models, ADELTA excels across various transpilation pairs, achieving pass@1 rate of 73.0 and 70.0 for PyTorch-Keras and PyTorch-MXNet transpilation pairs, respectively. These scores surpass those of the state-of-the-art LLM, GPT-4, by 17.4 and 15.0 points respectively.
- For training, we construct a PyTorch-Keras-MXNet corpus

*Published in the main track of IJCAI 2024.

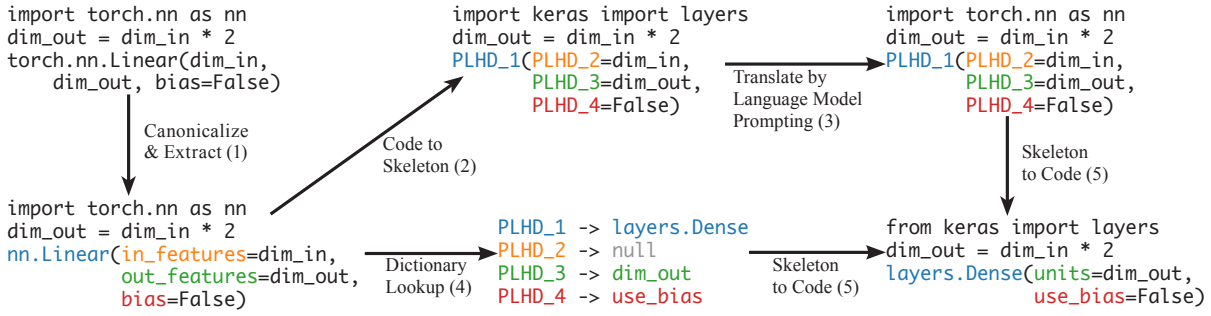


Figure 1: **An example of ADELt’s pipeline:** an import statement in the code skeleton is transpiled from PyTorch to Keras by a language model via few-shot prompting; a linear fully-connected layer is transpiled by removing the argument `in_features` and renaming other API keywords according to the learned dictionary. The number (1 to 5) near each arrow label corresponds to the step number in Section 2.

of deep learning code from various Internet sources, containing 49,705 PyTorch modules, 11,443 Keras layers/models, and 4,785 MXNet layers/models. We then build an evaluation benchmark for PyTorch-Keras and PyTorch-MXNet transpilation. The benchmark evaluates both our API keyword mapping algorithm and the overall source-to-source transpilation.

2 Method

ADELt (Adversarial DEep Learning Transpiler) is an algorithm that transpiles code from a source deep learning framework into an equivalent one in a target framework, by transpiling the skeletal code using a pretrained large language model, and then looking up each keyword in a dictionary learned with unsupervised domain-adversarial training. ADELt applies the following steps to each piece of input code, which we illustrate using the example shown in Figure 1:

1. Extract *API calls* from the source code. Such API calls can be automatically extracted with the Python’s built-in `ast` library. We then convert each API call into its canonical form, where each layer/function has a unique name, and all of its arguments are converted to keyword arguments. Finally, we extract all *API keywords* from the canonicalized API call, where an *API keyword* is the name of a layer/function or the name of a keyword argument.
2. Transform the program into its *code skeleton* by replacing each API keyword occurrence with a distinct placeholder.
3. Transpile the code skeleton, where all API keywords are replaced by placeholders, into the target DL framework using a pretrained big LM (e.g., Codex).
4. Look up each API keyword in the *API keyword dictionary*, and replace each keyword with its translation. To generate the API keyword dictionary, we first learn the API embeddings using domain-adversarial training based on contextual embeddings extracted by PyBERT (a BERT pretrained on Python code and then fine-tuned on deep learning code). Next, we calculate the cosine similarity between the embedding vectors. Then we generate the API keyword dictionary using a hierarchical algorithm.
5. Put each API keyword back into the transpiled code skeleton to generate the final output.

We describe each of these steps next in detail.

2.1 Canonicalization & API Keyword Extraction

We first parse the source code into an *abstract syntax tree* (AST) with the Python `ast` module. Then, canonicalization and API call extraction are applied to the AST.

Canonicalization. We canonicalize each API call using the following steps during both domain-adversarial training (Section 2.3) and inference. Each step involves a recursive AST traversal.

1. Unify the different import aliases of each module into the most commonly used name in the training dataset. For example, `torch.nn` is converted to `nn`.
2. Unify different aliases of each layer/function in a DL library into the name in which it was defined. We detect and resolve each alias by looking at its `__name__` attribute, which stores the callable’s original name in its definition.¹ For example, `layers.MaxPool2D` is converted to `layers.MaxPooling2D`.
3. Convert each positional argument of an API call into its equivalent keyword argument. Sort all keyword arguments according to the order defined in the function signature. This is done by linking the arguments of each API call to the parameters of its API signature using the `bind` method from Python’s `inspect` module.²

API keyword extraction. We define *API keyword* as the name of a layer/function or the name of a keyword argument. Once the input code is canonicalized, we locate each API keyword in the AST and then unparse the AST into the canonicalized source code.

2.2 Skeletal Code Transpilation

After canonicalizing the source program, ADELt then replaces all API keywords with a placeholder, turning the source program into its *code skeleton*. Each placeholder has textual form `PLACEHOLDER_i`, where $i = 1, 2, 3, \dots$. The code skeleton is then translated by Codex using few-shot prompting. The full prompt for this step is shown in Appendix A.4.

¹<https://docs.python.org/3/reference/datamodel.html#the-stand-ard-type-hierarchy>

²<https://docs.python.org/3/library/inspect.html#inspect.Signature.bind>

Algorithm 1 Pseudo-code for domain-adversarial training.

```

1 for (x1, y1), (x2, y2) in loader:
2   # N samples from X1, X2 respectively
3   # y1, y2: API keyword ids
4
5   h1 = B(x1).detach() # contextual embedding
6   h2 = B(x2).detach() # no gradient to PyBERT
7   z1 = G(h1) # generator hidden states
8   z2 = G(h2) # z1, z2: N x d
9
10  # dot product of z1 and output embeddings
11  logits1 = mm(z1, E1.view(d, m1))
12  logits2 = mm(z2, E2.view(d, m2))
13  L_CE1 = CrossEntropyLoss(logits1, y1)
14  L_CE2 = CrossEntropyLoss(logits2, y2)
15
16  # discriminator predictions
17  pred1 = D(z1)
18  pred2 = D(z2)
19  labels = cat(zeros(N), ones(N))
20  L_D = CrossEntropyLoss(pred1, labels)
21  L_G = CrossEntropyLoss(pred2, 1 - labels)
22
23  # joint update of G and E1
24  # to minimize L_CE1
25  optimize(G + E1 + E2, L_CE1 + L_CE2)
26  optimize(D, L_D) # train the discriminator
27  optimize(G, L_G) # train the generator

```

B: PyBERT used as the contextual embedder. G, D: the generator $\mathcal{G}$ and the discriminator $\mathcal{D}$.

E₁: a d by m_1 matrix, where the i -th column vector is the output embedding of API keyword $w_i^{(1)}$.

mm: matrix multiplication; cat: concatenation

2.3 Domain-Adversarial Training

Once the code skeleton is transpiled, we then transpile API keywords. We train the aligned embeddings of API keywords in a domain-adversarial setting. In Section 2.4, the embeddings will be used to generate a dictionary that maps an API keyword of the source deep learning framework $\mathcal{X}^{(1)}$ to an API keyword in the target DL framework $\mathcal{X}^{(2)}$.

Figure 2 illustrates the domain-adversarial approach of ADELt, and Algorithm 1 shows the pseudocode. A generator maps the contextual representations extracted by PyBERT into hidden states (line 5-8). The alignment of hidden states from different DL frameworks is enforced by the adversarial loss induced by the discriminator (line 17-21), so that output embeddings learned with these hidden states (line 11-14) are also aligned. Next, we describe each step in detail:

Each training example is a pair of API keyword occurrences with their context in the training corpus, denoted by $(x^{(1)}, x^{(2)})$. Each keyword occurrence $x^{(l)}$ is tokenized and encoded as multiple *byte pair encoding (BPE)* [Sennrich *et al.*, 2016] tokens. In our unsupervised setting, $x^{(1)}$ and $x^{(2)}$ are independent samples from $\mathcal{X}^{(1)}$ and $\mathcal{X}^{(2)}$ in the training dataset, respectively, and they are not necessarily translations of each other.

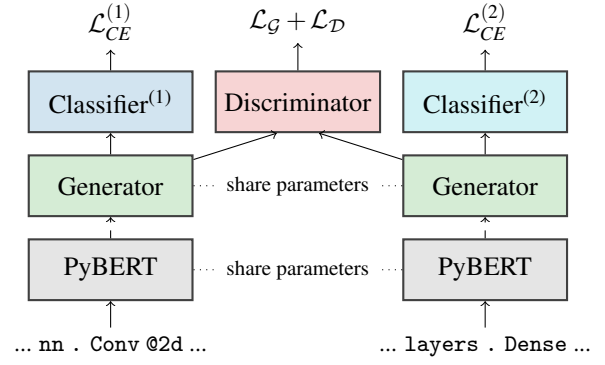


Figure 2: ADELt’s domain-adversarial training with contextual embeddings from a PyBERT. The generator and the PyBERT are shared between different DL frameworks. We do not fine-tune the PyBERT during adversarial training.

$$\begin{aligned}
\mathcal{L}_D &= -\mathbb{E}_{\text{data}}[\log \Pr_D(\text{pred} = 1 | \mathcal{G}(\mathbf{h}^{(1)}))] \\
&\quad - \mathbb{E}_{\text{data}}[\log \Pr_D(\text{pred} = 2 | \mathcal{G}(\mathbf{h}^{(2)}))] \\
\mathcal{L}_G &= -\mathbb{E}_{\text{data}}[\log \Pr_D(\text{pred} = 2 | \mathcal{G}(\mathbf{h}^{(1)}))] \\
&\quad - \mathbb{E}_{\text{data}}[\log \Pr_D(\text{pred} = 1 | \mathcal{G}(\mathbf{h}^{(2)}))]
\end{aligned} \tag{1}$$

$$\mathcal{L}_{CE}^{(l)} = -\mathbb{E}_{(x,y) \sim \text{data}^{(l)}} \left[\log \frac{\exp(\mathbf{z} \cdot \mathbf{e}_y^{(l)})}{\sum_{k=1}^{m^{(l)}} \exp(\mathbf{z} \cdot \mathbf{e}_k^{(l)})} \right] \tag{2}$$

PyBERT. *PyBERT* is our pretrained Transformer [Vaswani *et al.*, 2017; Devlin *et al.*, 2019] for Python code [Feng *et al.*, 2020; Kanade *et al.*, 2020; Roziere *et al.*, 2021]. Given a sequence of BPE tokens that represent an API keyword with its context $x^{(l)}$, PyBERT outputs a sequence of vectors—one vector in $\mathbb{R}^{d_b}$ for each token, where d_b is the hidden dimension size of PyBERT. We average-pool all BPE tokens of the keyword and get a single d_b -dimensional vector as the contextual embedding $\text{PyBERT}(x^{(l)})$ of the API keyword. We denote the contextual embedding of $x^{(1)}, x^{(2)}$ by $\mathbf{h}^{(1)}, \mathbf{h}^{(2)}$ respectively.

Generator and discriminator. We define two multi-layer perceptrons, a generator and a discriminator. A generator $\mathcal{G}$ encodes the contextual embeddings $\mathbf{h}^{(1)}, \mathbf{h}^{(2)}$ into hidden states $\mathbf{z}^{(1)}, \mathbf{z}^{(2)} \in \mathbb{R}^d$, and a discriminator $\mathcal{D}$ is trained to discriminate between $\mathbf{z}^{(1)}$ and $\mathbf{z}^{(2)}$. The generator is trained to prevent the discriminator from making accurate predictions, by making $\mathcal{G}(\text{PyBERT}(\mathcal{X}^{(1)}))$ and $\mathcal{G}(\text{PyBERT}(\mathcal{X}^{(2)}))$ as similar as possible. Our approach is inspired by domain-adversarial training [Ganin *et al.*, 2016], where domain-agnostic representations of images or documents are learned for domain adaptation. In our case, a domain is represented by a DL framework.

Formally, we define the probability $\Pr_D(\text{pred} = l | \mathbf{z})$ that a hidden state $\mathbf{z}$ is from the DL framework l predicted by the discriminator. Note that $\mathbf{z}^{(1)} = \mathcal{G}(\mathbf{h}^{(1)})$ and $\mathbf{z}^{(2)} = \mathcal{G}(\mathbf{h}^{(2)})$. The discriminator loss and the generator loss are computed as

the binary cross entropy against the true label and the reversed label, respectively, as shown in Equation (1).

Output embeddings. Our goal is to learn an embedding for each API keyword, but the contextual embedding of each keyword occurrence varies with its context. So we instead train a d -dimensional vector $\mathbf{e}_i^{(l)}$ for each API keyword $w_i^{(l)}$, such that $\mathbf{e}_i^{(l)}$ is similar to the generator hidden states $\mathbf{z}_j^{(l)}$ of this keyword’s occurrences and dissimilar to the hidden states $\mathbf{z}_k^{(l)}$ of any other keyword’s occurrences. $\mathbf{e}_i^{(l)}$ is considered the *output embedding* of the API keyword $w_i^{(l)}$. With similarity computed using dot product, our optimization objective is shown in Equation (2), equivalent to the cross-entropy loss of $m^{(l)}$ -way softmax-based classification.

Adversarial training. During each training iteration, the generator and discriminator are trained successively to minimize $\mathcal{L}_G$ and $\mathcal{L}_D$ respectively with mini-batch stochastic gradient descent. Minimizing the adversarial loss equals to minimizing the distance between two distributions of hidden states [Goodfellow *et al.*, 2014]. Therefore, the API keywords from the different DL frameworks will be mapped to an aligned embedding space.

Also, we jointly update the generator and the output embeddings to minimize $\mathcal{L}_{CE}^{(l)}$ with mini-batch SGD. The joint optimization is crucial, as updating the generator to minimize $\mathcal{L}_{CE}^{(l)}$ ensures that each generator hidden state $\mathbf{z}^{(l)}$ preserves enough information to recover its original API keyword. As a result, the output embeddings $\{\mathbf{e}_i^{(1)}\}_{i=1}^{m^{(1)}}$ and $\{\mathbf{e}_j^{(2)}\}_{j=1}^{m^{(2)}}$ are also aligned, as they are trained with vectors $\mathbf{z}^{(l)}$ from the aligned embedding space.

We do not fine-tune PyBERT during domain-adversarial training, as fine-tuning PyBERT makes the generator disproportionately strong that results in training divergence.

2.4 Hierarchical API Dictionary Generation

ADELTA calculates a *scoring matrix* using the aligned API keyword embeddings trained in Section 2.3. The entry in the i -th row and the j -th column of the matrix is the cosine similarity between $w_i^{(1)}$ and $w_j^{(2)}$, denoted by $s_{i,j}$. Given the scoring matrix, we need to generate an API keyword dictionary that maps each API keyword in one deep learning framework to an API keyword in another DL framework.

Greedy match is used to generate a dictionary in word translation of natural languages [Conneau *et al.*, 2018], where each source word is matched to the target word with the highest similarity score.

Structure of API keywords. Unlike words in NL, API keywords are *structured*: API keywords can be classified into two types based on their associated AST node: *callable names* (names of functions or classes), and *parameter names* (names of keyword arguments). In dictionary generation, we do not allow callable names to be translated to parameter names. We only allow parameter names to be translated to callable names in a special case when the weight passes a threshold. In this case, this parameter will be dropped and generate a new API call (the last case in Table 2). Another structural property is

that the matching of parameters depends on the matching of callables.

Hierarchical API dictionary generation algorithm leverages the structure of API keywords to generate a dictionary: **Step 1.** Consider each callable and its parameters as a group and compute the *group similarity* between each pair of groups, by summing up similarity scores in the greedy matching of parameter names, plus the similarity between two callable names. **Step 2.** Match groups greedily based on group similarity scores calculated in step 1.

3 Experiments

We evaluate the effectiveness of ADELTA on the task of transpilation between PyTorch, Keras, and MXNet³ and compare our method with baselines.

3.1 Skeletal Code Transpilation

We use Codex [Chen *et al.*, 2021], a LLM trained on public GitHub code, to transpile code skeletons. As an autoregressive language model trained on massive web data, Codex can handle translation tasks via prompting with few-shot demonstrations. Our prompt design aligns with Codex’s code translation setup, comprising a single input-output example and three instructions to keep placeholders unchanged. Appendix A.4 provides further details on this.

3.2 Training Setup

DL corpus. We consider 3 data sources **GitHub**, **JuiCe** [Agashe *et al.*, 2019], **Kaggle** [Quaranta *et al.*, 2021] to build our DL corpus:

- **GitHub:** The GitHub public dataset available on Google BigQuery.⁴ We keep py and ipynb files that contain torch, keras, or mxnet in the main and master branch of the repository. (69GB of clean Python code before filtering, 2.5GB after filtering)
- **JuiCe:** A code generation dataset [Agashe *et al.*, 2019] based on ipynb files from GitHub. JuiCe contains many files absent in the public dataset on Google BigQuery, since the latter is a selected subset of GitHub (10.7GB of clean Python code)
- **Kaggle:** All files in KGTorrent [Quaranta *et al.*, 2021], a dataset of Jupyter Notebooks from Kaggle.⁵ (22.1GB of clean Python code)

We tokenize all Python source code and extract subclasses of `torch.nn.Module`, `keras.layers.Layer`, or `keras.Model`. Then, we canonicalize (section 2.1) the code of each class definition. We byte-pair encode [Sennrich *et al.*, 2016], merge, and deduplicate codes from all sources. Finally, we collect all files into our *DL Corpus* containing

³We tried to evaluate using JAX [Bradbury *et al.*, 2018]. Sadly, JAX is a new DL framework and the GitHub corpus on BigQuery (based on a historical snapshot of GitHub) contains very few (318) examples of JAX.

⁴<https://console.cloud.google.com/marketplace/details/github/gi>thub-repos

⁵<https://kaggle.com>

	PyTorch-Keras						PyTorch-MXNet					
	F1		EM		Pass@1		F1		EM		Pass@1	
GPT-3 [Brown <i>et al.</i> , 2020]	26.6	32.0	22.4	26.0	23.4	27.2	25.8	32.8	23.4	25.0	25.0	26.4
Codex [Chen <i>et al.</i> , 2021]	59.9	67.1	51.5	54.6	53.4	57.6	57.4	69.0	53.2	56.2	54.2	57.6
GPT-4	67.7	74.9	55.6	64.6	56.8	66.0	60.3	71.8	54.0	60.2	55.0	60.8
Edit Distance (Cased)	31.2	30.1	20.3	16.8	20.3	16.8	37.7	35.7	22.8	21.0	22.8	21.0
Edit Distance (Uncased)	23.9	30.1	12.6	16.8	12.6	16.8	30.8	36.0	18.4	20.0	18.4	20.0
ADELTA (Small)	79.0	76.7	70.8	67.6	70.8	67.6	76.7	70.6	66.6	63.0	66.6	63.0
ADELTA (Base)	83.4	79.3	73.0	71.6	73.0	71.6	80.0	72.1	70.0	63.8	70.0	63.8

Table 1: **Comparison between ADELTA and other methods** on source-to-source transpilation. “ADELTA (Small)” is ADELTA with PyBERT_{SMALL} and “ADELTA (Base)” is ADELTA with PyBERT_{BASE}. There are two numbers in each table cell: the first one is for transpiling PyTorch to the other framework (Keras or MXNet), and the second one is for transpiling the other framework to PyTorch. Each number is the average of 5 runs with different random seeds.

49,705 PyTorch modules, 11,443 Keras layers/models, and 4,785 MXNet layers/models.

PyBERT is our Transformer encoder pretrained with the masked language modeling (MLM) [Devlin *et al.*, 2019] objective on all open-source Python files from the GitHub dataset. We consider two model sizes: PyBERT_{SMALL} (6-layer, 512-d) and PyBERT_{BASE} (12-layer, 768-d). Detailed pretraining hyperparameters are described in appendix A.1.

Adversarial training. The generator and discriminator of ADELTA are multilayer perceptrons. We search the learning rate and batch size according to the unsupervised validation criterion “*average cosine similarity*” [Conneau *et al.*, 2018], which measures the consistency between learned API keyword embeddings and generated keyword translations. Other hyperparameters are set based on previous studies [Conneau *et al.*, 2018] with details described in Appendix A.2.

3.3 Evaluation Benchmark

Our method is evaluated through the task of transpiling code snippets from one DL framework to another. Our benchmark consists of two parts: first, we use heuristics to identify potential matching pairs in the corpus, which were then refined through manual curation to ensure a solid evaluation benchmark; the second part of the benchmark includes a set of expert-transpiled examples, each accompanied by unit tests. For detailed methodology and statistics, please refer to Appendix A.3.

We report results in three evaluation metrics:

- **F1 score** quantifies the overlap between the predicted and ground truth outputs. In this context, we treat each prediction or ground truth as a bag of function calls. For each test case, we determine the number of exactly matched calls n_{match} , predicted calls n_{pred} , and ground truth calls n_{truth} . We define the F1 score for a particular example as $2n_{\text{match}}/(n_{\text{pred}} + n_{\text{truth}})$, and report the average F1 scores across all test cases.
- **Exact Match (EM) score** is a more rigorous metric that evaluates whether a model’s transpilation is exactly equivalent to the ground truth for each code snippet. It’s calculated as the proportion of exact matches to the total number of examples in the eval set.

- **Pass@1** assesses the proportion of examples for which the first transpilation attempt by the model successfully passes all the unit tests. These unit tests, created by experts for each benchmark example, evaluate the correctness of transpilation by execution.

3.4 Evaluation of Skeletal Code Transpilation

Transpiling code skeletons of DL programs is an easy task, and Codex easily learned transpilation patterns via few-shot prompting. In our evaluation benchmark, the exact match score of skeletal code transpilation using Codex is 100%.

3.5 Comparison with Other Methods

We compare ADELTA using PyBERT_{SMALL} and ADELTA using PyBERT_{BASE} with the following baselines. We run all methods 5 times with random seeds [10, 20, 30, 40, 50], and report the arithmetic average of all metrics.

End-to-end language models. We compare ADELTA with end-to-end few-shot LLM baselines, including GPT-3, Codex, and GPT-4, where the entire piece of source code, instead of the code skeleton, is fed into the LLM to generate the transpiled target program. For source-to-source translation, we use the “*completion*” endpoint of code-davinci-002 version for Codex and the “*chat*” endpoint of gpt-4-0314 for GPT-4. In both cases, we give the LLM a natural language instruction and 5 examples as demonstrations. Details of the prompts are shown in Appendix A.5.

Edit distance. We consider a rule-based baseline where we use edit distance [Levenshtein, 1966] as the similarity measure between API keywords, in place of the similarity measures calculated from learned embeddings. We apply hierarchical API dictionary generation exactly as what we do in ADELTA. We report the result of both cased and uncased setups for edit distance calculation.

The result is shown in Table 1. **ADELTA consistently outperforms other methods with respect to all metrics**, and it benefits from a larger pretrained PyBERT embedder. Importantly, ADELTA maintains its lead over GPT-4, even when GPT-4 is provided with additional examples for few-shot supervision. Moreover, ADELTA is much faster than GPT-4, because it uses a smaller LLM for transpiling the code skeleton, coupled with the dictionary lookup step, which requires only a tiny fraction of the time needed for LLM inference.

3.6 Case Studies

Table 2 shows four examples of PyTorch-Keras transpilation together with hypotheses of Codex and ADELt (Base). Both Codex and ADELt transpile the `nn.Conv2d` to Keras correctly by dropping the first argument `in_channels`. ADELt does not translate the parameter names of `nn.Embedding` to `input_dim` and `output_dim` correctly, while Codex does. However, we notice that Codex sometimes relies on the argument ordering heuristic. In the example of `nn.MultiheadAttention`, where parameters have a different ordering in Keras than in PyTorch, Codex generates the wrong translation, but ADELt successfully constructs the correct mapping between parameters.

Also, in the `nn.Embedding` example, Codex continues to generate code about “positional embeddings” after finishing transpilation. The extra code generated by Codex is relevant to the context.⁶ Still, the extra code should not be part of the translation. We have tried various ways to make Codex follow our instructions (see Appendix A.5 for details). However, because Codex is an end-to-end neural language model, our means of changing its predictions are limited, and the result is highly indeterministic. In the end, Codex still occasionally generates extra arguments or unneeded statements.

On the other hand, we decouple neural network training from the transpilation algorithm. ADELt transpiles between deep learning frameworks using deterministic keyword substitution based on a learned API keyword dictionary. The transpiled code is always syntactically correct. If a mistake is found in the dictionary (e.g., the `nn.Embedding` example in Table 2), it can be corrected by simply modifying the dictionary.

Correcting the API keyword dictionary by humans requires much less effort than building the dictionary manually from scratch, as ADELt generates a high-quality dictionary. Developers can even add additional rules to the transpiler. The flexibility of our decoupled design makes ADELt far easier to be integrated into real-world products than end-to-end neural translators/LMs are.

The last case in Table 2 shows an example where an API call (`layers.Dense` with `activation="relu"`) should be transpiled to two calls (`nn.Linear` and `nn.ReLU`). One-to-many mapping is rare in transpilation between deep learning frameworks, but the capability to model such mapping reflects the generality of a transpiler to other APIs. Both ADELt and Codex fail to solve this example because this usage is rarely seen in the training data. Still, if we train ADELt on an additional synthetic dataset (“ADELt+” in Table 2. See Appendix A.8 for details), it successfully solves this case, showing that our method can model one-to-many mappings when enough training data is available.

3.7 Ablation Studies

We conduct ablation studies on PyTorch-Keras transpilation to validate the contribution of each part of ADELt. As illustrated

⁶The definition of positional embeddings usually follows the definition of word embeddings (`nn.Embedding(vocab_size, ...)`) in the source code of a Transformer model.

in Figure 1, ADELt consists of five steps. Steps 1 (canonicalization and extraction), 2 (code to skeleton), and 5 (skeleton to code) are deterministic and always yield 100% accuracy as they do not rely on machine learning models. The accuracy of Step 3 (code skeleton transpilation) has been discussed in Section 3.4. Therefore, this section focus on the primary error source: the *API keyword translation* in Step 4. This crucial step involves mapping API keywords between frameworks. To evaluate its accuracy, we construct a high-quality dictionary by manually translating the top 50 most frequent API keywords from PyTorch to Keras. We then evaluate the translation’s efficacy using standard metrics: *precision@k* (for $k = 1, 5$) and the *mean reciprocal rank* (MRR) of correct translations. The results are shown in Table 3, where we study the following variants of ADELt:

Necessity of contextual embeddings. In “w/o PyBERT”, we replace PyBERT with Word2Vec [Mikolov *et al.*, 2013] embeddings of the same dimensions d_b trained on the same corpora. The result in Table 3 shows that this change significantly harms the performance of ADELt. This justifies the use of PyBERT, a high-quality pretrained representation of API keywords that can capture their contexts.

Contribution of adversarial loss. In “w/o Adv Loss”, we remove the adversarial loss during training. Instead, we only train the generator and the output embeddings with the cross-entropy loss in Equation (2). The result in Table 3 shows that adversarial training contributes ~ 6 pts in source-to-source transpilation, showing the effectiveness of adversarial training.

Comparison of similarity measures. By default, ADELt uses cosine similarity as the similarity measure for API dictionary generation. Table 3 shows the results of using dot product (inner). Measures based on cosine similarity outperforms dot product by a small margin. This fact implies that the performance of ADELt is insensitive to the choice of similarity measure.

4 Related Work

Source-to-source transpilation. Classical source-to-source transpilers use supervised learning. Nguyen *et al.* [2013] and Karaivanov *et al.* [2014] develop Java-C# transpilers using parallel corpora of open-source code. The dependency on parallel corpora renders these methods inapplicable to transpilation between deep learning frameworks, as parallel corpora are difficult to get.

Drawing inspiration from unsupervised neural machine translation (NMT) [Artetxe *et al.*, 2018], recent advancements have made unsupervised programming language translation possible [Lachaux *et al.*, 2020]. Such approaches, however, require vast amounts of in-domain unlabeled corpora, as evidenced by Lachaux *et al.* [2020] and Roziere *et al.* [2022], who utilized 744GB of GitHub source code and a dataset of 333k curated Java functions respectively. The scarcity of online deep learning code hinders their effectiveness for transpilation between DL frameworks, as we illustrate in Section 3.5.

Language models are few shot learners. GPT-3 [Brown *et al.*, 2020] is a language model with 175B parameters trained on massive web crawl data. GPT-3 can be applied to many

Source	<code>nn.Conv2d(64, 128, 3)</code>	Source	<code>nn.Embedding(vocab_size, embed_dim)</code>
Truth	<code>layers.Conv2D(filters=128, kernel_size=3)</code>	Truth	<code>layers.Embedding(input_dim=vocab_size, output_dim=embed_dim)</code>
Codex ✓	<code>layers.Conv2D(128, 3)</code>	Codex ✓	<code>layers.Embedding(vocab_size, embed_dim)</code> <code>self.position_emb = layers.Embedding(...)</code>
ADELTA ✓	<code>layers.Conv2D(filters=128, kernel_size=3)</code>	ADELTA ✗	<code>layers.Embedding(embeddings_initializer=embed_dim)</code>

Source	<code>nn.MultiheadAttention(model_dim, num_heads=num_heads, dropout=attn_dropout)</code>	Source	<code>in_dim = 256</code> <code>out_dim = 512</code> <code>layers.Dense(out_dim, activation='relu')</code>
Truth	<code>layers.MultiHeadAttention(num_heads=num_heads, key_dim=model_dim, dropout=attn_dropout)</code>	Truth	<code>in_dim = 256</code> <code>out_dim = 512</code> <code>nn.Linear(in_dim, out_dim)</code> <code>nn.ReLU()</code>
Codex ✗	<code>layers.MultiHeadAttention(model_dim, num_heads, dropout=attn_dropout)</code>	Codex ✗	<code>in_dim = 256</code> <code>out_dim = 512</code> <code>nn.Linear(in_dim, out_dim)</code>
ADELTA ✓	<code>layers.MultiHeadAttention(num_heads=num_heads, key_dim=model_dim, dropout=attn_dropout)</code>	ADELTA ✗	<code>in_dim = 256</code> <code>out_dim = 512</code> <code>nn.Linear(in_features=in_dim, out_features=out_dim)</code>
		ADELTA + ✓	<code>in_dim = 256</code> <code>out_dim = 512</code> <code>nn.Linear(in_features=in_dim, out_features=out_dim)</code> <code>nn.ReLU()</code>

Table 2: **Examples from the evaluation dataset of the PyTorch-Keras transpilation task and the Keras-PyTorch transpilation task.** We show the source code, ground truth target code, and the outputs from Codex, ADELTA, and ADELTA +. ✓: the output is the same or equivalent to the ground truth. ✓: the output contains an equivalent of the ground truth, but it also contains incorrect extra code. ✗: the output is incorrect.

	Keyword				Source Code	
	P@1	MRR			F1	
ADELTA (Small)	82.9	90.0	87.0	94.0	79.0	76.7
ADELTA (Base)	87.1	90.0	89.7	94.0	83.4	79.3
<i>Domain-adversarial training</i>						
w/o PyBERT (Small)	52.1	63.6	60.5	72.8	37.2	43.0
w/o PyBERT (Base)	45.0	54.6	56.8	66.0	33.0	36.3
w/o Adv Loss (Small)	80.4	88.6	85.3	93.1	65.8	73.6
w/o Adv Loss (Base)	86.3	90.5	89.3	94.3	78.2	72.3
<i>Measure for dictionary generation</i>						
Inner Product (Small)	81.3	79.6	86.3	85.4	74.6	73.2
Inner Product (Base)	85.4	93.2	88.8	95.7	80.2	78.8

Table 3: **Ablation study results.** By default, ADELTA is trained with the adversarial loss on contextual embeddings extracted by PyBERT, and then a dictionary is generated based on cosine similarity scores. We change one component of ADELTA (Small) or ADELTA (Base) in each experiment to assess its contribution.

NLP tasks without any gradient updates or fine-tuning, with tasks and few-shot demonstrations specified purely via text interaction with the model. Codex [Chen *et al.*, 2021] is a GPT-3 fine-tuned on publicly available code from GitHub, specialized for code generation tasks. GPT-4 is a LLM proficient in both code and NL trained using instruction finetuning. In contrast, the code generation step of ADELTA is keyword substitution instead of autoregressive generation. ADELTA outperforms

GPT-3, Codex, and GPT-4 in PyTorch-Keras transpilation and PyTorch-MXNet transpilation.

Adversarial learning & cross-lingual word embedding. Conneau *et al.* [2018] uses domain-adversarial [Ganin *et al.*, 2016] approach to align the distribution of two word embeddings, enabling natural language word translation without parallel data. The domain-adversarial training in ADELTA is inspired by their approach, but we align the distributions of the hidden states of *keyword occurrences* instead of API keyword embeddings.

5 Conclusion

In this work, we present ADELTA, an novel code transpilation algorithm for deep learning frameworks. ADELTA decouples code transpilation into two distinct phases: code skeleton transpilation and API keyword mapping. It leverages large language models (LLMs) for the transpilation of skeletal code, while using domain-adversarial training for the creation of an API keyword mapping dictionary. This strategic decoupling harnesses the strengths of two different model types. Through comprehensive evaluations using our specially curated PyTorch-Keras and PyTorch-MXNet benchmarks, we show that ADELTA significantly surpasses existing state-of-the-art transpilers in performance.

Acknowledgements

This work is supported in part by the U.S. National Science Foundation through grants IIS-1955488, IIS-2027575, ARO W911NF2110339, ONR N00014-21-1-2724, and DOE awards DE-SC0016260, DE-SC0021982.

References

- Rajas Agashe, Srinivasan Iyer, and Luke Zettlemoyer. Juice: A large scale distantly supervised dataset for open domain context-based code generation. *arXiv:1910.02216 [cs]*, Oct 2019. arXiv: 1910.02216.
- Mikel Artetxe, Gorka Labaka, Eneko Agirre, and Kyunghyun Cho. Unsupervised neural machine translation. *arXiv:1710.11041 [cs]*, Feb 2018. arXiv: 1710.11041.
- James Bradbury, Roy Frostig, Peter Hawkins, Matthew James Johnson, Chris Leary, Dougal Maclaurin, George Necula, Adam Paszke, Jake VanderPlas, Skye Wanderman-Milne, and Qiao Zhang. JAX: composable transformations of Python+NumPy programs, 2018.
- Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners. *arXiv:2005.14165 [cs]*, Jul 2020. arXiv: 2005.14165.
- Mark Chen, Jerry Tworek, Heewoo Jun, Qiming Yuan, Henrique Ponde de Oliveira Pinto, Jared Kaplan, Harri Edwards, Yuri Burda, Nicholas Joseph, Greg Brockman, Alex Ray, Raul Puri, Gretchen Krueger, Michael Petrov, Heidy Khlaaf, Girish Sastry, Pamela Mishkin, Brooke Chan, Scott Gray, Nick Ryder, Mikhail Pavlov, Alethea Power, Lukasz Kaiser, Mohammad Bavarian, Clemens Winter, Philippe Tillet, Felipe Petroski Such, Dave Cummings, Matthias Plappert, Fotios Chantzis, Elizabeth Barnes, Ariel Herbert-Voss, William Hebgen Guss, Alex Nichol, Alex Paino, Nikolas Tezak, Jie Tang, Igor Babuschkin, Suchir Balaji, Shantanu Jain, William Saunders, Christopher Hesse, Andrew N. Carr, Jan Leike, Josh Achiam, Vedant Misra, Evan Morikawa, Alec Radford, Matthew Knight, Miles Brundage, Mira Murati, Katie Mayer, Peter Welinder, Bob McGrew, Dario Amodei, Sam McCandlish, Ilya Sutskever, and Wojciech Zaremba. Evaluating large language models trained on code. *arXiv:2107.03374 [cs]*, Jul 2021. arXiv: 2107.03374.
- Alexis Conneau, Guillaume Lample, Marc’Aurelio Ranzato, Ludovic Denoyer, and Hervé Jégou. Word translation without parallel data. *arXiv:1710.04087 [cs]*, Jan 2018. arXiv: 1710.04087.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv:1810.04805 [cs]*, May 2019. arXiv: 1810.04805.
- Georgiana Dinu, Angeliki Lazaridou, and Marco Baroni. Improving zero-shot learning by mitigating the hubness problem. *arXiv:1412.6568 [cs]*, Apr 2015. arXiv: 1412.6568.
- Zhangyin Feng, Daya Guo, Duyu Tang, Nan Duan, Xiaocheng Feng, Ming Gong, Linjun Shou, Bing Qin, Ting Liu, Daxin Jiang, and Ming Zhou. Codebert: A pre-trained model for programming and natural languages. *arXiv:2002.08155 [cs]*, Sep 2020. arXiv: 2002.08155.
- Yaroslav Ganin, Evgeniya Ustinova, Hana Ajakan, Pascal Germain, Hugo Larochelle, François Laviolette, Mario Marchand, and Victor Lempitsky. Domain-adversarial training of neural networks. *arXiv:1505.07818 [cs, stat]*, May 2016. arXiv: 1505.07818.
- Ian J. Goodfellow, Jean Pouget-Abadie, Mehdi Mirza, Bing Xu, David Warde-Farley, Sherjil Ozair, Aaron Courville, and Yoshua Bengio. Generative adversarial networks. *arXiv:1406.2661 [cs, stat]*, Jun 2014. arXiv: 1406.2661.
- Herve Jegou, Cordelia Schmid, Hedi Harzallah, and Jakob Verbeek. Accurate image search using the contextual dissimilarity measure. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 32(1):2–11, Jan 2010.
- Aditya Kanade, Petros Maniatis, Gogul Balakrishnan, and Kensen Shi. Learning and evaluating contextual embedding of source code. *arXiv:2001.00059 [cs]*, Aug 2020. arXiv: 2001.00059.
- Svetoslav Karaivanov, Veselin Raychev, and Martin Vechev. Phrase-based statistical translation of programming languages. In *Proceedings of the 2014 ACM International Symposium on New Ideas, New Paradigms, and Reflections on Programming & Software*, Onward! 2014, page 173–184. Association for Computing Machinery, Oct 2014.
- Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv:1412.6980 [cs]*, Jan 2017. arXiv: 1412.6980.
- Marie-Anne Lachaux, Baptiste Roziere, Lowik Chanasot, and Guillaume Lample. Unsupervised translation of programming languages. *arXiv:2006.03511 [cs]*, Sep 2020. arXiv: 2006.03511.
- V. I. Levenshtein. Binary codes capable of correcting deletions, insertions and reversals. *Soviet Physics Doklady*, 10:707, Feb 1966. ADS Bibcode: 1966SPhD...10.707L.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. Roberta: A robustly optimized bert pretraining approach. Jul 2019.
- Tomas Mikolov, Kai Chen, Greg Corrado, and Jeffrey Dean. Efficient estimation of word representations in vector space. Jan 2013.
- Anh Tuan Nguyen, Tung Thanh Nguyen, and Tien N. Nguyen. Lexical statistical machine translation for language migration. In *Proceedings of the 2013 9th Joint Meeting on Foundations of Software Engineering*, ESEC/FSE 2013, page 651–654. Association for Computing Machinery, Aug 2013.
- Luigi Quaranta, Fabio Calefato, and Filippo Lanubile. Kgtorrent: A dataset of python jupyter notebooks from kaggle.

- 2021 *IEEE/ACM 18th International Conference on Mining Software Repositories (MSR)*, page 550–554, May 2021. arXiv: 2103.10558.
- Miloš Radovanović, Alexandros Nanopoulos, and Mirjana Ivanović. Hubs in space: Popular nearest neighbors in high-dimensional data. *Journal of Machine Learning Research*, 11(86):2487–2531, 2010.
- Baptiste Roziere, Marie-Anne Lachaux, Marc Szafraniec, and Guillaume Lample. Dobf: A deobfuscation pre-training objective for programming languages. *arXiv:2102.07492 [cs]*, Oct 2021. arXiv: 2102.07492.
- Baptiste Roziere, Jie M. Zhang, Francois Charton, Mark Harman, Gabriel Synnaeve, and Guillaume Lample. Leveraging automated unit tests for unsupervised code translation. *arXiv:2110.06773 [cs]*, Feb 2022. arXiv: 2110.06773.
- Rico Sennrich, Barry Haddow, and Alexandra Birch. Neural machine translation of rare words with subword units. *arXiv:1508.07909 [cs]*, Jun 2016. arXiv: 1508.07909.
- Ilya Sutskever, Oriol Vinyals, and Quoc V. Le. Sequence to sequence learning with neural networks. *arXiv:1409.3215 [cs]*, Dec 2014. arXiv: 1409.3215.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. Jun 2017.

A Appendix

A.1 PyBERT Pre-Training Hyperparameters and Implementation Details

The models are pre-trained with the RoBERTa [Liu *et al.*, 2019] pipeline in fairseq⁷ codebase. We pre-train each PyBERT on the GitHub dataset. On a NVIDIA DGX-2, it takes 8.2 hours and 23.1 hours to train PyBERT_{SMALL} and PyBERT_{BASE}, respectively.

Table 4 shows the pre-training hyperparameters of PyBERT_{SMALL} and PyBERT_{BASE}. We first pre-train each model on the Github dataset and then fine-tune it on our canonicalized PyTorch-Keras corpus. The learning rate is decayed according to the inverse square root schedule. We do not use early stopping — we use the last PyBERT checkpoint in ADELt.

Hyperparameter	PyBERT _{SMALL}	PyBERT _{BASE}
Number of layers	6	12
Hidden size d_b	512	768
FFN inner hidden size	2048	3072
Attention heads	8	12
Attention head size	64	64
Dropout	0.1	0.1
Attention dropout	0.0	0.0
FFN dropout	0.0	0.0
Adam β_1	0.9	0.9
Adam β_2	0.98	0.98
Adam ϵ	1e-6	1e-6
Weight decay	0.01	0.01
Gradient clipping	-	-
Peak learning rate	5e-4	5e-4
Batch size	2,048	2,048
Warmup steps	10,000	10,000
Total steps	125,000	125,000

Table 4: Pre-training hyperparameters of PyBERT

A.2 Domain-Adversarial Training Hyperparameters

The generator and the discriminator of ADELt are multilayer perceptrons. The activation function is ReLU for the generator and Leaky-ReLU for the discriminator. Dropout and label smoothing are applied for regularization. We train our generator, discriminator, and API keyword embeddings with Adam [Kingma and Ba, 2017] on 1,536,000 samples. There is a linear learning rate warmup over the first 10% of steps, and then we set the LR according to the invert square root decay rule. The learning rate scheduler uses linear warmup and inverse sqrt decay. The maximum learning rate is searched from [2e-4, 5e-4, 1e-3], and the batch size is searched from [64, 128, 256]. The peak learning rate and the batch size are searched according to the unsupervised validation criterion “average cosine similarity” [Conneau *et al.*, 2018] of the generated dictionary, which quantifies the consistency between the learned API keyword embeddings and the generated keyword translations. We set other hyperparameters according to prior

works [Conneau *et al.*, 2018], shown in table 5 (top). The learning rates and the batch sizes selected in the hyperparameter search are shown in table 5 (bottom). The total number of training steps is “total samples” (1,536,000) divided by the searched batch size, which is 6,000 steps for ADELt (Small) and 12,000 steps for ADELt (Base).

Generator activation	ReLU
Generator hidden size	2,048
Generator layers	1
Discriminator hidden size	2,048
Discriminator layers	1
Discriminator activation	LeakyReLU
Discriminator LeakyReLU slope	0.2
Dropout	0.1
Label smoothing	0.2
Warmup step ratio	10%
Adam β_1	0.9
Adam β_2	0.999
Adam ϵ	1e-8
Weight decay	0.001
Discriminator iterations per step	1
Total samples	1,536,000
Peak learning rate (Small)	2e-4
Batch size (Small)	128
Peak learning rate (Base)	5e-4
Batch size (Base)	256

Table 5: The hyperparameters of domain-adversarial training

A.3 Evaluation Data Collection

Our evaluation uses a parallel corpus of 100 examples derived from two distinct methods:

Heuristically Identified Parallel Examples: We sourced open-source projects on GitHub that benchmark various deep learning frameworks by replicating identical neural network architectures across those frameworks. Our parallel pair identification process relied on a heuristic comparing Python class names. Criteria for selection included pairs of PyTorch modules and Keras models/layers that (a) possessed identical class names and (b) achieved a BLEU score above 65. Following heuristic selection, we refined these pairs through manual extraction of code segments containing deep learning API calls, resulting in a corpus of 50 parallel examples. Then, human experts manually transpile those 50 PyTorch examples to MXNet, resulting in a corpus of 50 parallel examples for evaluating PyTorch-MXNet transpilation.

Expert-Transpiled Examples: The second set of 50 examples was assembled by selecting PyTorch module definitions from GitHub repositories with more than 1,000 stars and asking human experts to convert them into the Keras framework. The resulting PyTorch-Keras pairs tend to be longer and more challenging.

A.4 Details of Skeletal Code Transpilation

Table 6 shows by example how we transpile skeletal codes using Codex few-shot prompting.

⁷<https://github.com/facebookresearch/fairseq>

1. Each API keyword in the canonicalized source program is replaced with an distinct placeholder, numbered from 1 to n (the number of API keywords). The program after this step is called the code skeleton of the source program.
2. We append the code skeleton to the natural language prompt, *# Translate from PyTorch to Keras*, and four input-output pairs. The first three input-output pairs prompt the model to keep placeholders unchanged during transpilation. Our experiments show that three input-output pairs are required for 100% skeletal code transpilation correctness. Also, Codex can generalize to an arbitrary number of placeholders even if only three is given. The last input-output pair is a real example of PyTorch-Keras skeletal code transpilation.
3. This entire piece of input is fed into Codex, and Codex will complete this input by generating tokens after *# Keras*. The output of Codex is considered as the code skeleton of the target program.
4. Each placeholder is replaced with the API keyword in the target DL framework, by querying each API keyword before replacement (step 1) in the API keyword dictionary learned with ADELTA.

If the number of placeholders in the source skeleton and the number of placeholders in Codex’s output do not match, it is considered a failed example in evaluation. However, in practice, the success rate of skeletal code transpilation is 100% in our experiments. We attribute that to the fact that skeletal code in DL programs, in comparison to arbitrary Python code, tend to be high structured with fairly predictable import statements, constructors, and how the different DL layers are constructed and connected to each other.

A.5 Evaluation Setup of LLMs

Following the practices in Brown *et al.* [2020] and Chen *et al.* [2021], we use the “*completion*” endpoint of GPT-3 or Codex and the “*chat*” endpoint of GPT-4 for transpilation. We input some text as a prompt with few-shot demonstrations, and the model will generate a completion that attempts to match the prompt. table 7 shows two examples illustrating how we leverage GPT-3 or Codex for our task. table 8 shows two examples illustrating how we leverage GPT-4 for our task.

For source-to-source transpilation, prompt engineering is straightforward. In the PyTorch-Keras transpilation example, we tell the model to “*# Translate from PyTorch to Keras*” and then give 5 demonstrations from our evaluation dataset. Next, we input a piece of source code and “*# Keras*” and let the model generate a code completion starting from the following line. For chat models like GPT-4, we formulate the prompt in Markdown format, and extract the contents first Markdown code block as the model’s output. To prevent answers from being leaked to the language model, we do not allow any demonstration to share common API functions with the current evaluation example.

Prompt engineering of API keyword translation is trickier because there are two types of keywords. We represent callable names by one line containing its textual representation, and we represent parameter names by two lines, where the first

line is the name of the callable that the parameter belongs to, and the second line is the name of the parameter. We give 10 demonstrations from our evaluation dataset.

Although GPT-3 and Codex have strong capabilities in generating code related to our prompt, we find that they sometimes fail to follow our instructions to transpile between deep learning frameworks. We discuss this problem in section 3.5. We try several approaches to mitigate this issue:

1. Use the *Instruct* version⁸ of GPT-3/Codex: *text-davinci-002* and *code-davinci-002*.
2. Add a prefix to the input prompt based on simple rules. For example, if the source code starts with *nn.* in PyTorch, add *layers.* to the prompt and let the model generate a code completion after it. This trick is applicable to two examples shown in table 7.
3. Mask the logits of tokens that usually leads to irrelevant generations. Specifically, we find that the model tends to generate irrelevant extra code after a line break or random comments. So we add a bias of -100 to the logits of the hash mark “*#*”. We also add a bias of -100 to the logits of the line break if the source code contains no line breaks.

We find that these measures significantly improve the performance of GPT-3 and Codex on deep learning transpilation. All results of GPT-3 and Codex reported in section 3.5 are from the LMs with all these tricks turned on.

Conversely, GPT-4 did not exhibit similar issues, and given that it does not support logits masking, we engaged it directly using the prompts in Table 8 on *gpt-4-0314* without additional alterations.

A.6 Cross-Domain Local Scaling (CSLS)

Cross-Domain Local Scaling (CSLS) is a similarity measure for creating a dictionary based on high-dimensional embeddings. CSLS was proposed by Conneau *et al.* [2018] for word translation between natural languages. Empirical results by Conneau *et al.* [2018] show that using a pairwise scoring matrix (e.g. cosine similarity, dot product) in dictionary generation suffers from the *hubness* problem [Radovanović *et al.*, 2010], which is detrimental to generating reliable matching pairs as some vectors, dubbed *hubs*, are the nearest neighbors to many other vectors according to s , while others (anti-hubs) are not nearest neighbors of any point. This problem is observed in various areas [Jegou *et al.*, 2010; Dinu *et al.*, 2015]. CSLS is proposed to mitigate the hubness problem.

We also conduct an experiment to verify the effectiveness of CSLS in API keyword translation between deep learning frameworks. Specifically, we denote by $\mathcal{N}_s^{(l)}(w)$ the *neighborhood* of API keyword w , a set consisting of K elements with the highest similarity scores with w in DL framework $\mathcal{X}^{(l)}$. We calculate the average similarity score of $w_i^{(1)}$ to its neighborhood in DL framework $\mathcal{X}^{(2)}$ and denote it by $r_i^{(2)}$. Likewise, we denote by $r_j^{(1)}$ the average similarity score of $w_j^{(2)}$ to its neighborhood in DL framework $\mathcal{X}^{(1)}$. Then we

⁸<https://help.openai.com/en/articles/5832130-what-s-changed-with-engine-names-and-best-practices>

define a new similarity measure CSLS of $w_i^{(1)}$ and $w_i^{(2)}$ by subtracting $r_i^{(2)}$ and $r_j^{(1)}$ from their (doubled) similarity score $s_{i,j}$, as shown in eq. (3).

$$\begin{aligned} r_i^{(2)} &= \frac{1}{K} \sum_{k \in \mathcal{N}_s^{(2)}(w_i^{(1)})} s_{i,k} \\ r_j^{(1)} &= \frac{1}{K} \sum_{k \in \mathcal{N}_s^{(1)}(w_j^{(2)})} s_{k,j} \\ \text{CSLS}_{i,j} &= 2s_{i,j} - r_i^{(2)} - r_j^{(1)} \end{aligned} \quad (3)$$

CSLS can be induced from a parameter K and any similarity measure, including dot product and cosine similarity. Intuitively, compared with the score matrix of similarity measure s , the score matrix of CSLS assigns higher scores associated with isolated keyword pairs and lower scores of keywords lying in dense areas.

Given the (cosine similarity) scoring matrix scaled by CSLS, we then apply the hierarchical dictionary generation algorithm (section 2.4) to generate the API keyword dictionary. We search K in $\{5, 10, 20\}$ according to the unsupervised evaluation metric, and the result is similar, where $K = 5$ gives a slightly better result. table 9 shows the result of cosine-CSLS compared with cosine similarity.

table 9 shows that replacing cosine similarity with cosine-CSLS-5 does not impact the F1 score of transpiling PyTorch to Keras significantly, but it hurts the F1 score of transpiling Keras to PyTorch. The reason is that the vocabulary of API keywords is smaller than a natural language vocabulary. Hubness is not a problem for generating API keyword dictionaries; instead, penalizing the top-K may hurt the performance when there are relatively few valid candidates (e.g. Keras-to-PyTorch transpilation). Therefore, we do not use CSLS for ADELt.

A.7 Full Results with Error Bars

table 10 shows full results with error bars for PyTorch-Keras API keyword translation and source-to-source transpilation. The table includes the results of both the main comparison with GPT-3/Codex and ablation studies. We also add the results of GPT-3 and Codex on API keyword translation, where we randomly give the GPT-3 and Codex 10 examples as demonstrations. Details about prompt designs and hyperparameter setup are shown in appendix A.5. We do not calculate precision@5 and mean reciprocal rank for GPT-3 and Codex because the API provided by OpenAI does not support ranking a large number of generations cost-efficiently.

A.8 ADELt +

We created a new model, ADELt +, which is based on ADELt but trained on a synthetic dataset. Our goal is to evaluate whether our method can generalize to one-to-many mappings of APIs given enough data.

As we discussed in section 2.4, we allow parameter names to be translated to callable names when the weight passes a threshold τ . In this case, this parameter will be dropped and a new API call will be

generated. This mechanism allows ADELt to transpile layers.Dense(..., activation="relu") into two layers: nn.Linear(...) and nn.ReLU(), and similarly layers.Conv2D(..., activation="relu") into nn.Conv2D(...) and nn.ReLU(). However, such cases are rare in transpiling between deep learning frameworks, making it difficult to evaluate our model's ability to transpile one-to-many mappings in practice. Therefore, we create a synthetic dataset, where we replace all consecutive calls of layers.Dense and layers.ReLU in our dataset with layers.Dense(..., activation="relu"), and we replace all consecutive calls of layers.Conv2D and layers.ReLU with layers.Conv2D(..., activation="relu"). Then we train a new model, ADELt +, using our synthetic dataset.

We then evaluated ADELt + using our evaluation dataset. The value of the threshold τ is set heuristically to 5 (95% of values in the score matrix lies in -7 to 7). table 2 in section 3.6 and table 11 in the appendix show that ADELt + can model one-to-many mappings of APIs. For instance, table 11 shows that ADELt can transpile layers.Conv2D with activation='relu' into two API calls: nn.Conv2d and nn.ReLU.

A.9 More Case Studies

In Table 12 and Table 13, we select two PyTorch-Keras cases in our evaluation dataset for illustration. They are examples of the average length of all evaluation examples in the evaluation set.

In each case, ADELt makes the correct transpilation. The only textual difference is that ADELt's transpilation only contains keyword arguments while the ground truth still contains positional arguments. However, because the prediction and the ground truth are the same after canonicalization, we consider each case as an exact match during evaluation.

A.10 Deep Learning Transpilation across Different Programming Languages

In the main paper, all experiments are conducted on Python due to the scarcity of deep learning programs written in other programming languages such as Java or C. Despite that, in this section we show that ADELt is not limited to the same source and target languages by transpiling code written against the PyTorch library in Python 2 to Keras in Python 3.

To do so, we first canonicalize all PyTorch programs into Python 2 and all Keras programs into Python 3. Then we run ADELt on this modified training data to learn the API keyword dictionary. During inference, we transpile the code skeleton with Codex using the prompt shown in Table 14. Besides adding hint words such as "Python2" and "Python3" into the natural language prompt, we also find it necessary to add to the prompt some examples showing differences between Python 2 and Python 3, such as different print statements and different integer division operators. As is shown in Table 14, the skeletal codes were successfully transpiled from Python 2 and Python 3 along with the API keywords.

Canonicalized Source Program

```
import torch.nn as nn
dense = nn.Linear(in_features=dim_in, out_features=dim_out, bias=False)
```

Code Skeleton

```
import torch.nn as nn
dense = PLACEHOLDER_1(PLACEHOLDER_2=dim_in, PLACEHOLDER_3=dim_out, PLACEHOLDER_4=False)
```

Codex Input

Translate from PyTorch to Keras

PyTorch

PLACEHOLDER_1

Keras

PLACEHOLDER_1

PyTorch

PLACEHOLDER_2

Keras

PLACEHOLDER_2

PyTorch

PLACEHOLDER_3

Keras

PLACEHOLDER_3

PyTorch

```
import torch.nn as nn
```

```
class Model(nn.Module):
```

```
    def __init__(self):
```

```
        super().__init__()
```

```
        self.layer1 = PLACEHOLDER_1(PLACEHOLDER_2=16, PLACEHOLDER_3=32, PLACEHOLDER_4=3)
```

```
        self.layer2 = PLACEHOLDER_5()
```

```
    def forward(self, x):
```

```
        x = self.layer1(PLACEHOLDER_6=x)
```

```
        x = self.layer2(PLACEHOLDER_7=x)
```

```
        return x
```

Keras

```
import tensorflow.keras.layers as layers
```

```
class Model(layers.Layer):
```

```
    def __init__(self):
```

```
        super().__init__()
```

```
        self.layer1 = PLACEHOLDER_1(PLACEHOLDER_2=16, PLACEHOLDER_3=32, PLACEHOLDER_4=3)
```

```
        self.layer2 = PLACEHOLDER_5()
```

```
    def call(self, x):
```

```
        x = self.layer1(PLACEHOLDER_6=x)
```

```
        x = self.layer2(PLACEHOLDER_7=x)
```

```
        return x
```

PyTorch

```
import torch.nn as nn
```

```
dense = PLACEHOLDER_1(PLACEHOLDER_2=dim_in, PLACEHOLDER_3=dim_out, PLACEHOLDER_4=False)
```

Keras

Expected Codex Output

```
import tensorflow.keras.layers as layers
```

```
dense = PLACEHOLDER_1(PLACEHOLDER_2=dim_in, PLACEHOLDER_3=dim_out, PLACEHOLDER_4=False)
```

Target Program

```
import tensorflow.keras.layers as layers
```

```
dense = layers.Dense(units=dim_out, use_bias=False)
```

Table 6: Example inputs we give to Codex for skeletal code transpilation. We also show the expected outputs of the language model.

Source-to-Source Transpilation	Keyword Translation
<pre> # Translate PyTorch to Keras # PyTorch max_len = 512 self.embed_tokens = nn.Embedding(n_words, dim_emb) # Keras max_len = 512 self.embed_tokens = layers.Embedding(n_words, dim_emb, input_length=max_len) # PyTorch nn.Linear(dim_in, dim_out) # Keras layers.Dense(dim_out) (2 demonstrations omitted) # PyTorch F.log_softmax(logits, dim=-1) # Keras tf.nn.log_softmax(logits, axis=-1) # PyTorch nn.Conv2d(64, 128, 3) # Keras layers. </pre>	<pre> # Translate PyTorch to Keras # PyTorch F.log_softmax # Keras tf.nn.log_softmax # PyTorch nn.MaxPool2d stride # Keras layers.MaxPooling2D strides (7 demonstrations omitted) # PyTorch F.relu # Keras tf.nn.relu # PyTorch nn.Conv2d out_channels # Keras layers. </pre>
Conv2D(128, 3)	Conv2D filters

Table 7: Example inputs we give to GPT-3 or Codex for source-to-source transpilation and API keyword translation. We also show the expected outputs of the language models.

Source-to-Source Transpilation	Keyword Translation
You are an expert in deep learning. Transpile PyTorch to Keras: PyTorch: <pre> python max_len = 512 self.embed_tokens = nn.Embedding(n_words, dim_emb) </pre> Keras: <pre> python max_len = 512 self.embed_tokens = layers.Embedding(n_words, dim_emb, input_length=max_len) </pre> PyTorch: <pre> python nn.Linear(dim_in, dim_out) </pre> Keras: <pre> python layers.Dense(dim_out) </pre> <i>(2 demonstrations omitted)</i> PyTorch: <pre> python F.log_softmax(logits, dim=-1) </pre> Keras: <pre> python tf.nn.log_softmax(logits, axis=-1) </pre> PyTorch: <pre> python nn.Conv2d(64, 128, 3) </pre> Keras:	You are an expert in deep learning. Transpile PyTorch to Keras: PyTorch: <pre> python F.log_softmax </pre> Keras: <pre> python tf.nn.log_softmax </pre> PyTorch: <pre> python nn.MaxPool2d stride </pre> Keras: <pre> python layers.MaxPooling2D strides </pre> <i>(7 demonstrations omitted)</i> PyTorch: <pre> python F.relu </pre> Keras: <pre> python tf.nn.relu </pre> PyTorch: <pre> python nn.Conv2d out_channels </pre> Keras:
Conv2D(128, 3)	Conv2D filters

Table 8: Example inputs we give to GPT-4 for source-to-source transpilation and API keyword translation. We also show the expected outputs of the language models. Because GPT-4 outputs Markdown texts including both NL and code, we extract contents of the first Markdown code block as the output of the model.

	Keyword						Source-to-Source			
	P@1		P@5		MRR		BLEU		F1	
ADELTA (Small)	82.92	90.00	91.67	97.73	86.97	94.04	93.83	92.13	80.67	80.90
ADELTA (Base)	87.08	90.00	91.67	97.73	89.67	93.96	95.32	91.29	85.72	82.01
Inner Product (Small)	81.25	79.55	91.67	90.00	86.34	85.38	93.24	88.49	78.67	77.08
Inner Product (Base)	85.42	93.18	91.67	97.73	88.84	95.71	94.38	91.75	82.17	81.46
cos-CSLS-5 (Small)	84.17	83.18	97.92	93.64	89.89	89.12	94.24	90.43	83.17	76.60
cos-CSLS-5 (Base)	87.08	89.55	97.50	97.73	90.63	93.75	95.20	90.27	85.39	76.18

Table 9: **Results of CSLS.** By default, ADELTA computes similarity scores using *cosine similarity* to generate an API keyword dictionary. In this experiment, we replace cosine similarity with *inner product* or *cosine-CSLS-5* to compare different similarity measures. There are two numbers in each table cell: the first one is for transpiling PyTorch to PyTorch, and the second one is for transpiling Keras to PyTorch.

	Keyword		Source-to-Source	
	P@1		F1	
<i>LM few shot</i>				
GPT-3 [Brown <i>et al.</i> , 2020]	35.4 ± 6.1	39.1 ± 4.2	26.6 ± 5.1	32.1 ± 6.7
Codex [Chen <i>et al.</i> , 2021]	67.5 ± 8.3	79.1 ± 7.8	59.9 ± 2.7	67.1 ± 2.2
GPT-4	74.5 ± 5.8	83.2 ± 3.5	67.7 ± 2.6	74.9 ± 1.7
<i>ADELTA</i>				
ADELTA (Small)	82.9 ± 1.2	90.0 ± 1.6	79.0 ± 2.2	76.7 ± 1.5
ADELTA (Base)	87.1 ± 1.2	90.0 ± 2.5	83.4 ± 0.8	82.0 ± 2.2
w/o PyBERT (Small)	52.1 ± 2.6	63.6 ± 4.5	37.2 ± 8.2	43.0 ± 4.5
w/o PyBERT (Base)	45.0 ± 3.9	54.6 ± 5.3	33.0 ± 6.5	36.3 ± 3.2
w/o Adv Loss (Small)	80.4 ± 1.4	88.6 ± 2.0	65.8 ± 2.0	73.6 ± 1.8
w/o Adv Loss (Base)	86.3 ± 1.4	90.5 ± 2.4	78.2 ± 2.1	72.3 ± 3.6
Dot product (Small)	82.9 ± 4.5	90.0 ± 7.2	74.6 ± 2.7	73.2 ± 3.4
Dot product (Base)	87.1 ± 1.2	90.0 ± 2.0	80.2 ± 0.7	78.8 ± 0.8

Table 10: **Full results with 95% confidence intervals.** For each experiment, we run five experiments with different random seeds. Each cell has two intervals: the first one is for transpiling PyTorch to Keras, and the second one is for transpiling Keras to PyTorch. Each interval is the 95% confidence interval according to the Student’s t-Test, where we assume that the result of the five experiments follows a normal distribution.

Source	<pre> in_dim = 64 out_dim = 128 layers.Conv2D(filters=out_dim, kernel_size=3, activation="relu") </pre>
Truth	<pre> in_dim = 64 out_dim = 128 nn.Conv2d(in_dim, out_dim, 3) nn.ReLU() </pre>
Codex ✗	<pre> in_dim = 64 out_dim = 128 nn.Conv2d(in_dim, out_dim, 3) </pre>
ADELT ✗	<pre> in_dim = 64 out_dim = 128 nn.Linear(in_features=in_dim, out_features=out_dim, kernel_size=3) </pre>
ADELT + ✓	<pre> in_dim = 64 out_dim = 128 nn.Conv2d(in_channels=in_dim, out_channels=out_dim, kernel_size=3) nn.ReLU() </pre>

Table 11: **A synthetic example of convolution layer from the evaluation dataset of the Keras-PyTorch transpilation task.** We show the Keras code, ground truth PyTorch code, and the outputs from Codex, ADELT, and ADELT +. ✓: the output is the same or equivalent to the ground truth. ✓: the output contains an equivalent of the ground truth, but it also contains incorrect extra code. ✗: the output is incorrect.

```

# Source Program
import torch.nn as nn
class BasicBlock(nn.Module):
    def __init__(self, dim):
        super.__init__()
        self.bn1 = nn.BatchNorm2d(dim)
        self.act1 = nn.LeakyReLU(0.2)
        self.conv1 = nn.Conv2d(dim, dim, 3)
        self.pool1 = nn.MaxPool2d(3, 2)

# Transpiled by ADEL T
import tensorflow.keras.layers as layers
class BasicBlock(layers.Layer):
    def __init__(self, dim):
        super.__init__()
        self.bn1 = layers.BatchNormalization()
        self.act1 = layers.LeakyReLU(alpha=0.2)
        self.conv1 = layers.Conv2D(filters=dim, kernel_size=3)
        self.pool1 = layers.MaxPooling2D(pool_size=3, stride=2)

# Ground Truth
import tensorflow.keras.layers as layers
class BasicBlock(layers.Layer):
    def __init__(self, dim):
        super.__init__()
        self.bn1 = layers.BatchNormalization()
        self.act1 = layers.LeakyReLU(0.2)
        self.conv1 = layers.Conv2D(dim, 3)
        self.pool1 = layers.MaxPooling2D(3, 2)

```

Table 12: Additional case study 1.

```

# Source Program
import torch.nn as nn
class AttentionBlock(nn.Module):
    def __init__(self, args):
        super().__init__()
        self.attn = nn.MultiheadAttention(
            args.d_model, args.n_heads, dropout=args.att_dropout)
        self.drop1 = nn.Dropout(args.dropout)
        self.norm1 = nn.LayerNorm(args.d_model)

# Transpiled by ADEL T
import tensorflow.keras.layers as layers
class AttentionBlock(layers.Layer):
    def __init__(self, args):
        super().__init__()
        self.attn = layers.MultiHeadAttention(
            num_heads=args.n_heads, key_dim=args.d_model, dropout=args.att_dropout)
        self.drop1 = layers.Dropout(rate=args.dropout)
        self.norm1 = layers.LayerNormalization()

# Ground Truth
import tensorflow.keras.layers as layers
class AttentionBlock(layers.Layer):
    def __init__(self, args):
        super().__init__()
        self.attn = layers.MultiHeadAttention(
            args.n_heads, args.d_model, dropout=args.att_dropout)
        self.drop1 = layers.Dropout(args.dropout)
        self.norm1 = layers.LayerNormalization()

```

Table 13: Additional case study 2.

Canonicalized Source Program in Python 2

```
import torch.nn as nn
dense = nn.Linear(in_features=dim_in / 2, out_features=dim_out / 2, bias=False)
```

Code Skeleton

```
import torch.nn as nn
print dim_in, dim_out
dense = PLACEHOLDER_1(PLACEHOLDER_2=dim_in / 2, PLACEHOLDER_3=dim_out / 2, PLACEHOLDER_4=False)
```

Codex Input

Translate from PyTorch in Python2 to Keras in Python3

PyTorch in Python2

PLACEHOLDER_1

Keras in Python3

PLACEHOLDER_1

PyTorch in Python2

PLACEHOLDER_2

Keras in Python3

PLACEHOLDER_2

PyTorch in Python2

PLACEHOLDER_3

Keras in Python3

PLACEHOLDER_3

PyTorch in Python2

```
import torch.nn as nn
```

```
class Model(nn.Module):
```

```
    def __init__(self):
```

```
        super().__init__()
```

```
        print "Building Model"
```

```
        self.layer1 = PLACEHOLDER_1(PLACEHOLDER_2=16 / 2, PLACEHOLDER_3=32, PLACEHOLDER_4=3)
```

```
        self.layer2 = PLACEHOLDER_5()
```

```
    def forward(self, x):
```

```
        x = self.layer1(PLACEHOLDER_6=x)
```

```
        x = self.layer2(PLACEHOLDER_7=x)
```

```
        return x
```

Keras in Python3

```
import tensorflow.keras.layers as layers
```

```
class Model(layers.Layer):
```

```
    def __init__(self):
```

```
        super().__init__()
```

```
        print("Building Model")
```

```
        self.layer1 = PLACEHOLDER_1(PLACEHOLDER_2=16 // 2, PLACEHOLDER_3=32, PLACEHOLDER_4=3)
```

```
        self.layer2 = PLACEHOLDER_5()
```

```
    def call(self, x):
```

```
        x = self.layer1(PLACEHOLDER_6=x)
```

```
        x = self.layer2(PLACEHOLDER_7=x)
```

```
        return x
```

PyTorch in Python2

```
import torch.nn as nn
```

```
print dim_in, dim_out
```

```
dense = PLACEHOLDER_1(PLACEHOLDER_2=dim_in / 2, PLACEHOLDER_3=dim_out / 2, PLACEHOLDER_4=False)
```

Keras in Python3

Codex Output

```
import tensorflow.keras.layers as layers
```

```
print(dim_in, dim_out)
```

```
dense = PLACEHOLDER_1(PLACEHOLDER_2=dim_in // 2, PLACEHOLDER_3=dim_out // 2, PLACEHOLDER_4=False)
```

Target Program in Python 3

```
import tensorflow.keras.layers as layers
```

```
print(dim_in, dim_out)
```

```
dense = layers.Dense(units=dim_out // 2, use_bias=False)
```

Table 14: Example of transpiling from PyTorch in Python 2 to Keras in Python 3.