

# NEURON-LEVEL ANALYSIS OF CULTURAL UNDERSTANDING IN LARGE LANGUAGE MODELS

**Taisei Yamamoto**

The University of Tokyo, Riken  
yamamo96@is.s.u-tokyo.ac.jp

**Ryoma Kumon**

The University of Tokyo, Riken  
kumoryo9@is.s.u-tokyo.ac.jp

**Danushka Bollegala**

University of Liverpool  
danushka@liverpool.ac.uk

**Hitomi Yanaka**

The University of Tokyo, Riken  
hyanaka@is.s.u-tokyo.ac.jp

## ABSTRACT

As large language models (LLMs) are increasingly deployed worldwide, ensuring their fair and comprehensive cultural understanding is important. However, LLMs exhibit cultural bias and limited awareness of underrepresented cultures, while the mechanisms underlying their cultural understanding remain underexplored. To fill this gap, we conduct a neuron-level analysis to identify neurons that drive cultural behavior, introducing a gradient-based scoring method with additional filtering for precise refinement. We identify both *culture-general* neurons contributing to cultural understanding regardless of cultures, and *culture-specific* neurons tied to an individual culture. These neurons account for less than 1% of all neurons and are concentrated in shallow to middle MLP layers. We validate their role by showing that suppressing them substantially degrades performance on cultural benchmarks (by up to 30%), while performance on general natural language understanding (NLU) benchmarks remains largely unaffected. Moreover, we show that *culture-specific* neurons support knowledge of not only the target culture, but also related cultures. Finally, we demonstrate that training on NLU benchmarks can diminish models’ cultural understanding when we update modules containing many *culture-general* neurons. These findings provide insights into the internal mechanisms of LLMs and offer practical guidance for model training and engineering. Our code is available at <https://github.com/ynklab/CULNIG>

## 1 INTRODUCTION

LLMs are rapidly spreading throughout the world with their ability to solve various tasks. Our world is culturally diverse, and our knowledge, commonsense, and values are not always universal. LLMs must possess cultural understanding to be deployed fairly and prevent cultural inequity. However, several studies have pointed out that LLMs, which are mainly trained on English-dominant corpora, often exhibit culture-related biases, generating outputs skewed toward certain highly represented cultures (Naous et al., 2024; Myung et al., 2024; Sukiennik et al., 2025). In order to evaluate the cultural understanding of LLMs, a number of benchmarks have been constructed (Myung et al., 2024; Chiu et al., 2025; Rao et al., 2025; Zhao et al., 2024, *inter alia*). Additionally, some methods have been proposed to enhance cultural awareness of LLMs (Li et al., 2024a;b; Liu et al., 2025). Nonetheless, the mechanisms behind the cultural understanding of LLMs have not been well investigated. In order to improve the cultural understanding of LLMs efficiently and robustly, it is desirable to elucidate the inner workings by which LLMs perform culture-related inference.

Previous studies have applied neuron-level analysis to investigate various properties of LLMs, such as social bias (Yang et al., 2024) and personality (Deng et al., 2025). Regarding cultural mechanisms, Ying et al. (2025) analyzed neurons activated most strongly when the prompt language aligns with the cultural content. In addition, Namazifard & Galke (2025) proposed a method to disentangle culture neurons from language neurons. These studies primarily examine culture in relation to language, rather than the mechanisms by which LLMs shape their behavior based on cultural infor-

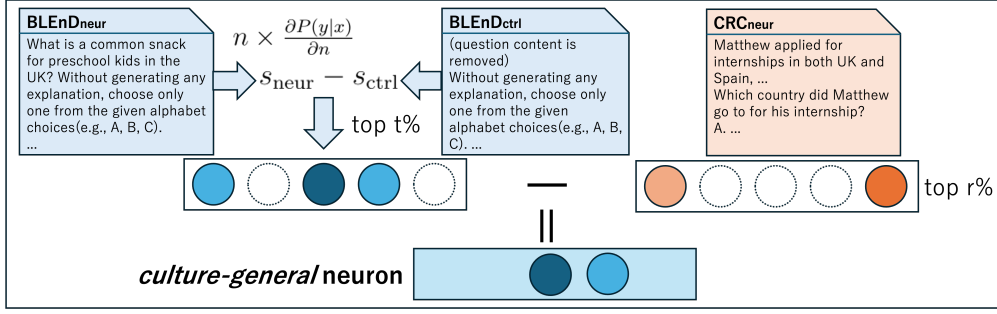


Figure 1: An overview of CULNIG when identifying *culture-general* neurons. We first select the top  $t\%$  of the neurons ranked by gradient-based attribution scores on  $\text{BLENDeur} - \text{BLENDectrl}$  ( $s_{\text{neur}} - s_{\text{ctrl}}$ ) to find neurons contributing to cultural mechanisms. By subtracting  $s_{\text{ctrl}}$ , we exclude neurons facilitating task understanding. We then remove the top  $r\%$  of the neurons on  $\text{CRCneur}$  to filter out superficial neurons activated by country names.

mation. Moreover, they rely on activation-based methods, which can be imprecise because cultural representations are not necessarily encoded in every token of culturally relevant texts.

In this paper, we explore three research questions: (i) the existence and distribution of *culture-general* neurons that contribute to cultural understanding across cultures, (ii) the differences of *culture-specific* neurons across cultures and the correlation between these neurons and cultural relations, and (iii) the potential engineering applications of our neuron analysis. We interpret cultural understanding along two dimensions: (a) knowledge specific to particular cultures and (b) the ability to capture differences in values across cultural backgrounds. To address these questions, we introduce **CULTure Neuron Identification Pipeline with Gradient-based Scoring (CULNIG, Figure 1)**, a method to accurately identify neurons that contribute to the cultural understanding of LLMs. CULNIG employs gradient-based attribution scores to rank neurons and a control dataset to exclude neurons associated with task understanding. We also construct the CountryRC (Country Reading Comprehension, CRC) dataset to filter out superficial neurons.

We comprehensively evaluate identified neurons using both cultural benchmarks and general natural language understanding (NLU) benchmarks that do not necessarily require cultural understanding. As a result, masking *culture-general* neurons significantly degrades the cultural understanding of LLMs while having only minor impacts on the performance in NLU benchmarks. Importantly, although CULNIG leverages problems from only a subset of cultural knowledge categories, the identified neurons generalize to broader cultural mechanisms, encompassing different knowledge domains, cultural values, and multilingual settings. *Culture-general* neurons account for fewer than 1% of all neurons and are concentrated in MLP modules of shallow to middle layers. We further show that masking *culture-specific* neurons leads to LLMs losing cultural knowledge of the target and related cultures. Moreover, we demonstrate that when we fine-tune a model with NLU datasets, updating modules containing many *culture-general* neurons can cause greater degradation of cultural understanding after training. These findings illustrate how insights into the inner workings of LLMs can inform practical engineering decisions.

## 2 RELATED WORK

### 2.1 EVALUATING CULTURAL UNDERSTANDING OF LLMs

Several cultural benchmarks have been developed to measure the cultural understanding of LLMs. BLENDe (Myung et al., 2024) covers everyday knowledge across 16 cultures in six categories, with multilingual short answer questions and English multiple-choice questions (MCQs). CulturalBench (Chiu et al., 2025) is an MCQ benchmark of cultural knowledge spanning 45 countries. NormAd (Rao et al., 2025) evaluates cultural etiquette through daily-life scenarios, asking whether the behaviors are acceptable in the target country. WorldValuesBench (Zhao et al., 2024), derived from World Values Survey (WVS) Wave 7 (Haerpfer et al., 2020), assesses understanding of cultural values by a prediction task of survey responses based on demographic attributes.

Prior studies have pointed out that LLMs often exhibit cultural biases toward highly represented cultures in training corpora (Naous et al., 2024; Myung et al., 2024; Sukiennik et al., 2025). Ying et al. (2025) demonstrates Cultural-Linguistic Synergy, a phenomenon where the performance of LLMs on cultural benchmarks improves when the prompt language agrees with the cultural content. In contrast, Myung et al. (2024) reports that Cultural-Linguistic Synergy does not always appear for low-resource languages, where limited language proficiency may act as a bottleneck. Building on these studies, we analyze the cultural understanding and behavior of LLMs at the neuron level, utilizing existing cultural benchmarks.

## 2.2 NEURON-BASED INTERPRETABILITY ANALYSIS

Mechanistic interpretability attempts to uncover the internal mechanisms of black-box LLMs, with many studies focusing on neurons as the unit of analysis. Dai et al. (2022) proposed a gradient-based attribution method to identify neurons that express a certain knowledge. They show that only a few knowledge neurons in deep layers support factual recall in BERT (Devlin et al., 2019). Using similar gradient-based attribution, Chen et al. (2025) located query-relevant neurons that facilitate question answering, and Yang et al. (2024) found bias neurons and mitigated bias by pruning them.

Moreover, many methods have been proposed to identify neurons based on their activation probability. Tang et al. (2024) and Kojima et al. (2024) identified language-specific neurons that are activated when LLMs are prompted in a specific language, with the former introducing LAPE (Language Activation Probability Entropy). Regarding cultural understanding, Ying et al. (2025) analyzed neurons underlying Culture-Linguistic Synergy. Namazifard & Galke (2025) proposed CAPE (Culture Activation Probability Entropy) to isolate culture neurons from language-specific neurons of LAPE, using a dataset of culturally diverse texts and entropy measures.

However, these methods often lack comprehensive evaluation across multiple cultural understanding benchmarks. Moreover, although both positive and negative activations encode useful information, activation-based approaches consider only positive activations, while clipping negative activations to zero activation probability. Thus, activation-based methods are typically limited to modules with nonlinear activation functions where negative values are clipped to zero. Also, since cultural content is not necessarily expressed in every token, unlike languages, activation probabilities may not be suitable for identifying culture neurons. Therefore, we adopt a gradient-based attribution approach and validate identified neurons across multiple benchmarks spanning different cultural attributes.

## 3 METHODS

In this section, we introduce CULNIG to identify *culture-general* and *culture-specific* neurons that directly support cultural understanding, respectively. Removing these neurons is expected to substantially alter model behavior on cultural benchmarks, unlike neurons that merely respond to culture-related tokens.

### 3.1 NEURONS IN LLMs

Each layer of a neural network can be represented as a hidden vector where its dimensions correspond to neurons. For the concept of neurons in an LLM, we follow Yu & Ananiadou (2024). In transformer-based LLMs (Vaswani et al., 2017), an input sequence  $[t_1, t_2, \dots, t_T]$  with  $T$  tokens is first mapped to token embeddings: the  $i$ -th token  $t_i$  is transformed into  $\mathbf{h}_i^{(0)} \in \mathbb{R}^d$ , where  $d$  is the hidden size. These embeddings are then processed by  $L$  layers, each consisting of an attention module and a multilayer perceptron (MLP). The  $l$ -th layer transforms its input  $\mathbf{h}_i^{(l-1)}$  as:

$$\mathbf{h}_i^{(l)} = \mathbf{h}_i^{(l-1)} + \mathbf{a}_i^{(l)} + \mathbf{f}_i^{(l)} \quad (1)$$

$\mathbf{a}_i^{(l)}$  and  $\mathbf{f}_i^{(l)}$  denote the outputs of the attention and MLP modules, respectively.

In (multi-head) attention layers, query, key, and value vectors are first computed as  $\mathbf{q}_i^{(l)} = \mathbf{W}_q^{(l)} \mathbf{h}_i^{(l-1)}$ ,  $\mathbf{k}_i^{(l)} = \mathbf{W}_k^{(l)} \mathbf{h}_i^{(l-1)}$ , and  $\mathbf{v}_i^{(l)} = \mathbf{W}_v^{(l)} \mathbf{h}_i^{(l-1)}$ , where  $\mathbf{W}_q^{(l)}, \mathbf{W}_k^{(l)}, \mathbf{W}_v^{(l)} \in \mathbb{R}^{DH \times d}$  denote query, key, and value matrices.  $D$  is the head dimension and  $H$  is the number of heads. Then, each vector is

split into  $H$  heads, and the outputs for each head  $h$  are calculated as follows:

$$\alpha_i^{(l,h)} = \text{softmax}\left(\sqrt{\frac{1}{D}}(q_i^{(l,h)} \cdot k_1^{(l,h)}, \dots, q_i^{(l,h)} \cdot k_i^{(l,h)})\right) \quad (2)$$

$$o_i^{(l,h)} = \sum_{j=1}^i \alpha_{i,j}^{(l,h)} v_j^{(l,h)} \quad (3)$$

Head outputs are gathered with the output matrices  $\mathbf{W}_o^{(l,h)} \in \mathbb{R}^{d \times D}$  to obtain the final output as:

$$\mathbf{a}_i^{(l)} = \sum_{h=1}^H \mathbf{W}_o^{(l,h)} o_i^{(l,h)} \quad (4)$$

In MLP layers, recent LLMs commonly employ gated linear unit (Shazeer, 2020), expressed as:

$$\mathbf{f}_i^{(l)} = \mathbf{W}_{\text{down}}^{(l)} \left( \sigma(\mathbf{W}_{\text{gate}}^{(l)}(\mathbf{h}_i^{(l-1)} + \mathbf{a}_i^{(l)})) \odot \mathbf{W}_{\text{up}}^{(l)}(\mathbf{h}_i^{(l-1)} + \mathbf{a}_i^{(l)}) \right) \quad (5)$$

$\mathbf{W}_{\text{gate}}^{(l)}, \mathbf{W}_{\text{up}}^{(l)} \in \mathbb{R}^{N \times d}$ ,  $\mathbf{W}_{\text{down}}^{(l)} \in \mathbb{R}^{d \times N}$  are projection matrices,  $\sigma$  is the activation function, and  $N$  is the intermediate size.

Geva et al. (2021) show that MLP modules can be interpreted as key-value memories. The output  $\mathbf{f}_i^{(l)}$  is expressed as a weighted sum of the column vectors of  $\mathbf{W}_{\text{down}}^{(l)}$  (subvalues), and the weights are computed as the inner products of the inputs and the row vectors of  $\mathbf{W}_{\text{gate}}^{(l)}$  and  $\mathbf{W}_{\text{up}}^{(l)}$  (subkeys). The  $k$ -th neuron in the  $l$ -th layer gate projection is given by  $n_{\text{gate}}^{(l,k)} = (\mathbf{W}_{\text{gate}}^{(l)}(\mathbf{h}_i^{(l-1)} + \mathbf{a}_i^{(l)}))_k$ , which functions as a weight for the corresponding subvalue. By analyzing the contribution of these intermediate neurons, MLP outputs can be decomposed into a sum of subvalues. For MLPs, we focus on neurons in the gate projection, since the gate and up projections share the same subvalue, and gate neurons play a role as a gate to determine whether they pass the weights. Similarly, the output of an attention module can be decomposed into a weighted sum of the column vectors of  $\mathbf{W}_o^{(l,h)}$ , and the weights are determined by query, key, and value vectors. Thus, we search for neurons from the query, key, and value modules.

### 3.2 NEURON ATTRIBUTION SCORES

In order to quantify the importance of each neuron on a given instance, we adopt the method based on Yang et al. (2024). Let  $P(y|x)$  denote the probability of the output sequence  $y$  assigned by the model when given an input sequence  $x$ . The attribution score of the neuron  $n^{(l,k,i)}$  at the  $i$ -th token position is calculated using the following formula:

$$s^{(l,k,i)}(x, y) = n^{(l,k,i)} \times \frac{\partial P(y|x)}{\partial n^{(l,k,i)}} \quad (6)$$

We then take the maximum score across token positions:

$$s^{(l,k)}(x, y) = \max_i s^{(l,k,i)}(x, y) \quad (7)$$

Note that Equation 6 can be viewed as a first-order approximation of the causal effect of neuron  $n^{(l,k,i)}$ . Let  $P(y|x, n^{(l,k,i)} = u)$  denote the output probability when the activation value of  $n^{(l,k,i)}$  is  $u$ , and  $\bar{u}$  denote the actual activation. The causal effect of  $n^{(l,k,i)}$  on probability is  $P(y|x, n^{(l,k,i)} = \bar{u}) - P(y|x, n^{(l,k,i)} = 0)$ . Here, we expand  $P(y|x, n^{(l,k,i)} = u)$  around  $\bar{u}$  using the Taylor expansion as follows:

$$P(y|x, n^{(l,k,i)} = u) \approx P(y|x, n^{(l,k,i)} = \bar{u}) + \frac{\partial P(y|x, n^{(l,k,i)} = \bar{u})}{\partial \bar{u}} \times (u - \bar{u}) \quad (8)$$

When we set  $u = 0$ , we obtain the following formula:

$$s^{(l,k,i)}(x, y) = \bar{u} \times \frac{\partial P(y|x, n^{(l,k,i)} = \bar{u})}{\partial \bar{u}} \approx P(y|x, n^{(l,k,i)} = \bar{u}) - P(y|x, n^{(l,k,i)} = 0) \quad (9)$$

To calculate the causal effects of all neurons, we have to run the inference by masking each neuron one-at-a-time, which requires an enormous computational cost because LLMs typically contain millions of neurons (Table 12). In contrast, we can efficiently calculate  $s^{(l,k,i)}(x, y)$  in a single run.

We aggregate the score on a dataset  $D$  with  $Q$  instances as the weighted sum over the exact probability:

$$s^{(l,k)}(D) = \sum_{q=1}^Q P(y_q|x_q) \times s^{(l,k)}(x_q, y_q) \quad (10)$$

This is because when the model predicts the correct answer with higher confidence, it should contain more reliable information.

### 3.3 NEURON SELECTION

To identify *culture-general* and *culture-specific* neurons, we use the MCQs from BLEnD, which provide sufficient instances and reduce the risk of overfitting to individual examples. BLEnD covers 16 countries and six categories, ensuring diversity in cultural topics. To test whether identified neurons generalize across different domains of cultural knowledge, we split BLEnD by category: three categories (*food, work-life, sport*) for neuron identification (BLEnD<sub>neur</sub>) and the remaining three (*education, family, holidays/celebrations/leisure*) for evaluation (BLEnD<sub>test</sub>). BLEnD provides 500 questions, and each question has multiple instances derived from different answer choices. We sample up to five instances per question to balance the number of instances of each question, yielding 12,701 instances in BLEnD<sub>neur</sub> and 10,331 in BLEnD<sub>test</sub>.

Moreover, we prepare BLEnD<sub>ctrl</sub> to isolate neurons that contribute purely to cultural inference. In BLEnD<sub>ctrl</sub>, the question content is removed, leaving only the answer choices and the instruction for the answer format (Table 4). Neuron scores are calculated as  $s^{(l,k)}(\text{BLEnD}_{\text{neur}}) - s^{(l,k)}(\text{BLEnD}_{\text{ctrl}})$ , so that we can exclude neurons related to other properties, such as task understanding.

Table 1: An example of the CountryRC (CRC) dataset.

| Passage                                                                                                                                                           | Question                                                   |
|-------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------|
| Matthew applied for internships in both {country_A} and {country_B}, but only the company in {country_A} responded. He accepted and worked there over the summer. | Which country did Matthew go to for his internship? A. ... |

Since BLEnD evaluates culturally dependent knowledge, all the problems explicitly include country names. Thus, the top-scoring neurons on BLEnD<sub>neur</sub> may contain superficial neurons that simply respond to tokens of country names rather than cultural content. To filter out such superficial neurons, we construct another control dataset called CountryRC (CRC)<sup>1</sup>, in which the correct answer is always a country name that appears in the context (Table 1). We utilize ChatGPT<sup>2</sup> to create CRC. Answering CRC requires models to recognize and propagate information about the country name, but it does not involve cultural understanding. CRC contains 50 problems per country, with half used for neuron identification (CRC<sub>neur</sub>), and the remainder for evaluation (CRC<sub>test</sub>).

For *culture-general* neurons, we first select the top  $t\%$  of neurons ranked by  $s^{(l,k)}(\text{BLEnD}_{\text{neur}}) - s^{(l,k)}(\text{BLEnD}_{\text{ctrl}})$ , and then exclude the top  $r\%$  ranked by  $s^{(l,k)}(\text{CRC}_{\text{neur}})$ . This procedure defines CULNIG-general. For *culture-specific* neurons of a country  $c$ , we first apply the same process to select neurons using only the instances of  $c$  in the datasets, with an additional filtering step. Specifically, the score for  $c$  is calculated as  $s^{(l,k,c)} = s^{(l,k)}(\text{BLEnD}_{\text{neur}}^{(c)}) - s^{(l,k)}(\text{BLEnD}_{\text{ctrl}}^{(c)})$ . We compute the z-score of each neuron over the 16 countries in BLEnD as  $z^{(c)} = \frac{s^{(l,k,c)} - \mu}{\sigma}$ , where  $\mu$  and  $\sigma$  are the mean and standard deviation of  $s^{(l,k,c)}$  across countries. Neurons with  $z^{(c)} < 0.5$  are removed, as they are likely to contribute to multiple cultures. This threshold of z-score is determined through a preliminary experiment. The whole pipeline defines CULNIG-specific.

<sup>1</sup><https://huggingface.co/datasets/Taise228/CountryRC>

<sup>2</sup><https://chatgpt.com/overview>

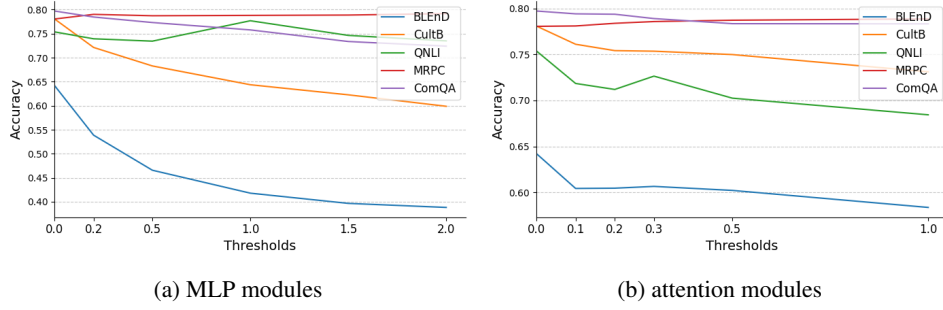


Figure 2: Accuracy of gemma-3-12b-it on each benchmark as more top-scoring neurons on BLENd (threshold  $t$ ) are masked, with neurons selected from MLP and attention modules, respectively.

## 4 EXPERIMENT AND ANALYSIS

In this section, we first describe our experimental settings in Section 4.1. We then compare the roles of each module and decide the thresholds in Section 4.2. Based on its result, we identify *culture-general* neurons in Section 4.3 and *culture-specific* neurons in Section 4.4. Finally, we show a potential application of our findings from an engineering perspective in Section 4.5.

### 4.1 MODELS AND DATASETS

In our experiments, we use gemma-3-12b-it, gemma-3-27b-it (Gemma Team, 2025), Qwen-3-14B (Qwen Team, 2025), Llama-3-8B-Instruct (Grattafiori et al., 2024), phi-4 (Abdin et al., 2024), and Falcon3-10B-Instruct (TII Team, 2024). We select various state-of-the-art open-source models to demonstrate the robustness and generalizability of our findings (see Appendix B for details).

As explained in Section 3.3, we use  $\text{BLENd}_{\text{neur}}$ ,  $\text{BLENd}_{\text{ctrl}}$ , and  $\text{CRC}_{\text{neur}}$  for neuron identification. For evaluation, we employ  $\text{BLENd}_{\text{test}}$  and CulturalBench (CultB) to measure cultural knowledge, NormAd as a task involving both cultural knowledge and values, and WorldValuesBench (WVB) to assess understanding of cultural values. We also use short answer questions (SAQs) of  $\text{BLENd}_{\text{test}}$  to evaluate LLMs in a different task and multilingual settings. In addition, we utilize four NLU benchmarks:  $\text{CRC}_{\text{test}}$ , CommonsenseQA (ComQA) (Talmor et al., 2019), QNLI, and MRPC (Wang et al., 2019), as comparison tasks that do not necessarily require cultural understanding.

Regarding evaluation metrics, we use accuracy (%) for all benchmarks except WVB. For WVB, we frame the task as a prediction of a questionnaire response given the country. The questionnaire uses a Likert scale, and we adopt the  $\text{score}_c$  metric based on Xu et al. (2025):

$$\text{score}_c = \frac{1}{N} \sum_{n=1}^N \left( 1 - \frac{|a_c^{(n)} - p_c^{(n)}|}{\text{max distance}} \right) \times 100 \quad (11)$$

$a_c^{(n)}$  is the majority answer among participants from country  $c$ ,  $p_c^{(n)}$  is the model prediction, and max distance is the maximum possible distance between the options and  $a_c^{(n)}$ . A higher  $\text{score}_c$  indicates greater alignment.

Considering the sensitivity of LLMs to task instructions (Zhan et al., 2024), we prepare four prompt formats for each benchmark using ChatGPT (for BLENd, the task instruction is included in the questions, so we prompt them without additional instructions). Further details are given in Appendix A.

### 4.2 ROLES OF MODULES: ATTENTION VS MLP

First, we conduct a preliminary experiment to analyze the roles of each module and decide the threshold in CULNIG-general. We separately select neurons from MLP and attention modules of gemma-3-12b-it, varying the threshold  $t$  for the top-ranked neurons. We fix the threshold for  $\text{CRC}_{\text{neur}}$  to  $r = 1\%$ . Figure 2 shows the evaluation results when masking the identified neurons.

Table 2: Evaluation results of masking *culture-general* (cult) and random (rand) neurons. Random scores are averaged over ten seeds of neuron selection. Values in parentheses denote standard deviations. **Bold** values indicate statistically significant score reductions relative to the random scores.

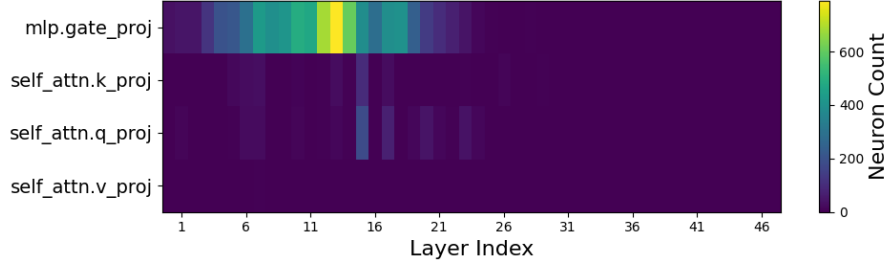
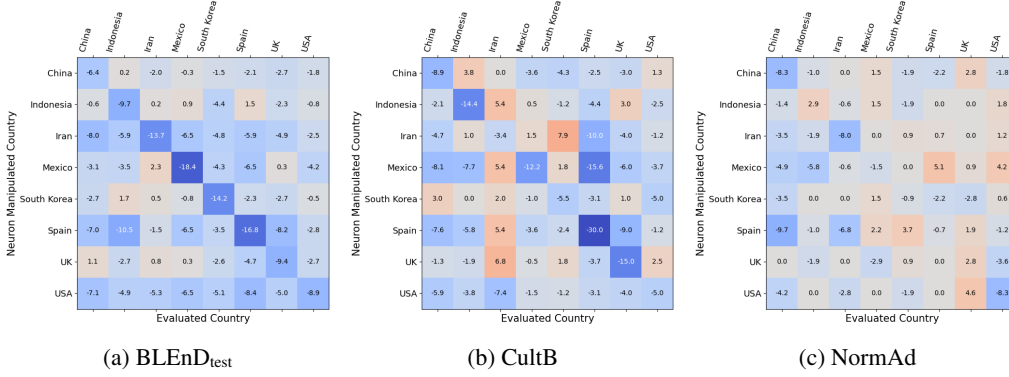
| Model                         | #Neuron | BLEND <sub>test</sub> | CultB        | NormAd       | WVB          | ComQA        | QNLI         | MRPC        |
|-------------------------------|---------|-----------------------|--------------|--------------|--------------|--------------|--------------|-------------|
| Chance rate                   | -       | 25.00                 | 25.00        | 33.33        | 49.85        | 20.00        | 50.00        | 50.00       |
| gemma-3<br>-12b-it            | orig    | 0                     | 64.22        | 78.08        | 58.54        | 64.08        | 79.71        | 75.37       |
|                               | cult    | 8,087                 | <b>37.93</b> | <b>62.00</b> | <b>52.02</b> | <b>58.46</b> | <b>75.10</b> | 72.77       |
|                               | rand    | 8,087                 | 63.57(0.46)  | 77.31(0.28)  | 57.55(0.57)  | 64.03(0.59)  | 79.18(0.60)  | 75.46(4.81) |
| gemma-3<br>-27b-it            | orig    | 0                     | 61.37        | 81.32        | 58.76        | 64.47        | 80.88        | 91.43       |
|                               | cult    | 14,273                | <b>39.96</b> | 69.76        | <b>52.31</b> | 60.98        | 79.32        | 90.81       |
|                               | rand    | 14,273                | 62.17(3.15)  | 78.32(7.11)  | 57.19(1.62)  | 62.07(7.03)  | 78.40(6.71)  | 87.56(9.87) |
| Qwen3<br>-14B                 | orig    | 0                     | 65.96        | 76.92        | 56.85        | 65.22        | 81.76        | 71.31       |
|                               | cult    | 7,340                 | <b>35.84</b> | <b>57.07</b> | <b>49.02</b> | <b>60.70</b> | <b>75.23</b> | 76.20       |
|                               | rand    | 7,340                 | 65.47(0.49)  | 75.98(0.40)  | 56.26(0.65)  | 64.46(1.04)  | 80.86(0.42)  | 71.49(1.2)  |
| Llama-<br>3.1-8B-<br>Instruct | orig    | 0                     | 60.18        | 70.54        | 47.71        | 64.05        | 76.74        | 64.43       |
|                               | cult    | 4,268                 | <b>32.19</b> | <b>36.94</b> | <b>37.65</b> | <b>51.68</b> | <b>51.97</b> | 48.64       |
|                               | rand    | 4,268                 | 57.75(0.97)  | 67.25(1.03)  | 43.88(1.59)  | 61.55(1.71)  | 72.84(1.24)  | 55.78(6.05) |
| phi-4                         | orig    | 0                     | 63.89        | 78.30        | 59.68        | 65.0         | 80.43        | 89.15       |
|                               | cult    | 7,447                 | <b>35.05</b> | <b>57.72</b> | 51.84        | 66.48        | <b>70.60</b> | 85.84       |
|                               | rand    | 7,447                 | 63.29(0.63)  | 76.94(1.71)  | 56.38(2.98)  | 61.82(2.67)  | 78.89(2.10)  | 86.98(1.93) |
| Falcon3<br>-10B-<br>Instruct  | orig    | 0                     | 57.98        | 71.74        | 55.26        | 58.00        | 79.73        | 74.57       |
|                               | cult    | 9,282                 | <b>35.47</b> | <b>56.81</b> | <b>48.75</b> | 59.16        | <b>71.85</b> | 70.30       |
|                               | rand    | 9,282                 | 57.64(0.31)  | 71.07(0.23)  | 54.06(1.39)  | 57.4(0.83)   | 78.89(0.71)  | 74.17(3.19) |

We find that masking MLP neurons causes substantial degradation on cultural benchmarks, while accuracies on QNLI and MRPC remain unaffected. For ComQA, the accuracy shows a moderate drop, likely because ComQA contains culture-related questions (e.g., What island country is ferret popular?  $\rightarrow$  great britain). Beyond  $t = 1\%$ , declines on BLEND<sub>test</sub> and CultB become gradual and parallel those on QNLI and ComQA, indicating that additional neurons contribute less specifically to cultural understanding. For attention neurons, the overall impact is smaller, but the scores on cultural benchmarks and QNLI decline to some extent. Although QNLI is solvable only with in-context information, it contains cultural sentences (e.g., What is the first major city in the stream of the Rhine?). Cultural knowledge can help solve QNLI, so the reduction may come from lost cultural understanding. Therefore, attention modules can still contain culture neurons. Beyond  $t = 0.2\%$ , the slopes on cultural benchmarks are similar to those on QNLI and ComQA.

These results corroborate prior studies showing that transformer MLPs primarily support knowledge recall, whereas attention modules facilitate in-context information processing (Meng et al., 2022; Ortu et al., 2024). This observation suggests that LLMs can rely more heavily on MLP neurons to solve cultural benchmarks, which require recall of out-of-context knowledge. Based on these observations, we adopt different thresholds for MLP and attention neurons in CULNIG-general, setting  $t_{\text{MLP}} = 1\%$  and  $t_{\text{attn}} = 0.2\%$ . In CULNIG-specific, we do not separate MLP and attention neurons, since the z-score-based filtering step can remove neurons that facilitate task understanding. We set  $t = 0.3\%$  and  $r = 1\%$ , reflecting the expectation that *culture-specific* neurons are fewer than *culture-general* neurons.

### 4.3 CULTURE-GENERAL NEURONS

With the settings described in Section 4.2, we identify *culture-general* neurons. Table 2 shows the evaluation results when suppressing *culture-general* neurons and random neurons averaged over ten seeds. In the table, p-values are defined as the probability that the score reduction with random neurons is greater than or equal to that with *culture-general* neurons, estimated by setting a bootstrapping sample size to 2,000. Here, the scores of random neurons are computed in two ways: the average over ten seeds and the score of a uniformly sampled single seed (for sensitivity analysis). If both p-values are smaller than 0.05, *culture-general* neurons are regarded as statistically significant.

Figure 3: The distribution of *culture-general* neurons in gemma-3-12b-it.Figure 4: Score reductions after masking *culture-specific* neurons of gemma-3-12b-it.

We observe that eliminating *culture-general* neurons consistently causes significant degradation on cultural benchmarks, while the impact on NLU benchmarks is smaller. In particular, for BLEND<sub>test</sub>, the score drops substantially up to 30%, although the identified neurons account for fewer than 1% of the total. For CRC<sub>test</sub>, the models achieved almost 100% accuracy both before and after masking neurons (Table 18), suggesting that few superficial neurons were included. Notably, although neurons are identified solely using specific cultural knowledge categories, performance also declines on the unseen categories (BLEND<sub>test</sub>), demonstrating generalization beyond knowledge domains. This generalization further extends across task formats (CultB) and even across cultural attributes, such as cultural etiquette and values (NormAd and WVB). Moreover, we evaluate the models on SAQs of BLEND<sub>test</sub> and demonstrate that masking *culture-general* neurons degrades the accuracy in the multilingual setting as well (Table 13, Appendix D). These results imply that *culture-general* neurons capture a broad representation of cultural understanding. Figure 3 shows the distribution of *culture-general* neurons in gemma-3-12b-it. Most of the neurons are located in shallow to middle MLP modules, and this tendency is consistent across models (Appendix C), suggesting that CULNIG-general captures a general property of LLMs. We also show the ablation studies of each step in CULNIG-general in Appendix H.

#### 4.4 CULTURE-SPECIFIC NEURONS

Next, we apply CULNIG-specific to identify *culture-specific* neurons that support understanding of individual cultures. We focus on eight countries covered in all of BLEND, CultB, and NormAd (China, Indonesia, Iran, Mexico, South Korea, Spain, UK, and USA), which are culturally diverse in the Inglehart-Welzel World Cultural Map from WVS Wave 7 (Haerpfer et al., 2020).

Figure 4 shows score reductions when masking *culture-specific* neurons in gemma-3-12b-it. For BLEND<sub>test</sub> and CultB, the largest drops occur in the target cultures, confirming that identified neurons are associated with knowledge of the target culture. Meanwhile, this pattern is less clear for NormAd, suggesting that cultural knowledge and other properties, such as etiquette and values, might be represented by different neurons. Moreover, *culture-specific* neurons tend to affect related cultures. For example, masking neurons corresponding to Mexico most strongly affects the prob-

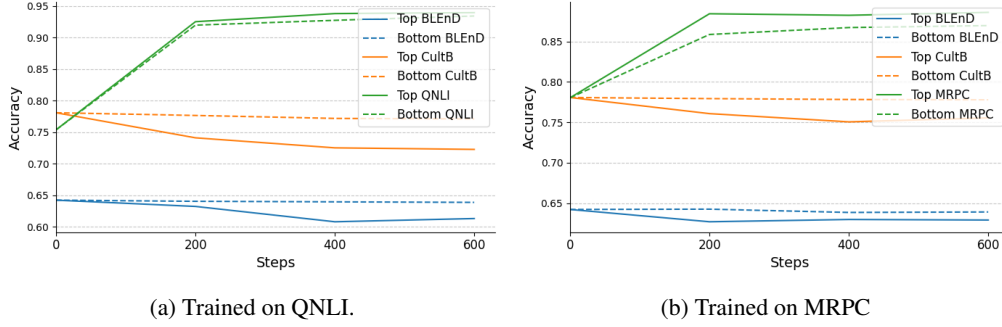


Figure 5: Evaluation results of gemma-3-12b-it on BLEnD<sub>test</sub>, CultB, QNLI, and MRPC when fine-tuned on QNLI or MRPC, updating only 10% of the total parameters. Updated modules are selected either from those containing the most *culture-general* neurons (Top) or those without *culture-general* neurons (Bottom).

lems of Mexico (the mean rank of score reduction among 16 cultures over six models is 1.17), and the second most affected culture is Spain (the mean rank was 3.83). We observe that historically or geographically related cultures tend to affect each other, indicating that the neurons underlying the related cultures are shared (Table 15).

The distribution of *culture-specific* neurons is similar to that of *culture-general* neurons (Figure 11a). In contrast, these results differ from CAPE, which reported that culture neurons are concentrated in the upper layers. We replicated the experiments of CAPE with gemma-3-12b-it, but could not reproduce it, failing to separate culture and language neurons. Consequently, LAPE and CAPE neurons had negligible impacts on evaluation scores. Further investigation of this discrepancy is left for future work. One possible factor is that we use gradient-based attribution scores, whereas CAPE uses activation-based scores. Another factor is that we evaluate LLMs based on accuracy on QA tasks, while CAPE uses the perplexity metric on cultural texts. A detailed account of this evaluation is presented in Appendix F.

#### 4.5 APPLICATIONS: TARGET MODULE SELECTION FOR TRAINING

In this section, we demonstrate a potential application of our findings from an engineering perspective. Fine-tuning LLMs often risks degrading their abilities on other tasks (Luo et al., 2025) and also requires enormous computational costs. To achieve robust and efficient training, we propose to select updating modules based on their roles.

We fine-tune (a language model of) gemma-3-12b-it with QNLI and MRPC, updating only a portion of the modules. For module selection, we sort the modules by the number of *culture-general* neurons, and select either those with the most *culture-general* neurons (top-culture modules) or those with none (bottom-culture modules) until the number of parameters exceeds 10%. When an MLP gate projection is selected, we also include the corresponding up and down projections, and when a query, key, or value module is selected, we also include the corresponding query, key, value, and out projections, since neurons in those modules are connected as subkeys and subvalues (Section 3.1). We fine-tune the model for 600 steps with a learning rate of 3e-5 and evaluate it every 200 steps.

The selected top-culture modules are all MLP modules from shallow to middle layers, while the bottom-culture modules mainly consist of very shallow attention modules and very deep attention and MLP modules (Table 17). The evaluation results are shown in Figure 5. We observe that the target scores (QNLI or MRPC) improved in both cases. However, when updating the top-culture modules, the scores of cultural benchmarks decrease. Meanwhile, updating the bottom-culture modules has little effect on cultural abilities. These results suggest that we can train the model efficiently and robustly by selecting target components based on their roles. The details and experiments with different parameter settings are shown in Appendix G.

## 5 CONCLUSION AND LIMITATIONS

We introduced CULNIG, a pipeline to identify neurons that contribute to the cultural understanding of LLMs. We evaluated six LLMs with *culture-general* neurons masked and demonstrated that the scores on the cultural benchmark decreased significantly, while the impacts on the NLU benchmarks were minor. Although identified with a limited domain of cultural knowledge problems, these neurons affected broader cultural attributes, including understanding of cultural values and performance on cultural knowledge benchmarks even in multilingual settings. Moreover, we located *culture-specific* neurons that are tied to individual cultures and confirmed that masking these neurons impaired knowledge of both the target and related cultures. *Culture-general* and *culture-specific* neurons were concentrated in shallow to middle MLP layers. Finally, we demonstrated that when we fine-tuned LLMs on NLU benchmarks, cultural understanding was more easily lost by updating modules containing many *culture-general* neurons than by updating modules without *culture-general* neurons. While our findings do not directly improve the cultural understanding of LLMs, they provide a foundation for future studies to do so.

## ACKNOWLEDGMENTS

This work was supported by JST CREST Grant Number JPMJCR2565, Japan.

## REFERENCES

- Marah Abdin, Jyoti Aneja, Harkirat Behl, Sébastien Bubeck, Ronen Eldan, Suriya Gunasekar, Michael Harrison, Russell J. Hewett, Mojan Javaheripi, Piero Kauffmann, James R. Lee, Yin Tat Lee, Yuanzhi Li, Weishung Liu, Caio C. T. Mendes, Anh Nguyen, Eric Price, Gustavo de Rosa, Olli Saarikivi, Adil Salim, Shital Shah, Xin Wang, Rachel Ward, Yue Wu, Dingli Yu, Cyril Zhang, and Yi Zhang. Phi-4 technical report. *arXiv preprint arXiv:2412.08905*, 2024. URL <https://arxiv.org/abs/2412.08905>.
- Lihu Chen, Adam Dejl, and Francesca Toni. Identifying query-relevant neurons in large language models for long-form texts. In *Proceedings of the Thirty-Ninth AAAI Conference on Artificial Intelligence and Thirty-Seventh Conference on Innovative Applications of Artificial Intelligence and Fifteenth Symposium on Educational Advances in Artificial Intelligence, AAAI’25/IAAI’25/EAAI’25*. AAAI Press, 2025. ISBN 978-1-57735-897-8. doi: 10.1609/aaai.v39i22.34529. URL <https://doi.org/10.1609/aaai.v39i22.34529>.
- Yu Ying Chiu, Liwei Jiang, Bill Yuchen Lin, Chan Young Park, Shuyue Stella Li, Sahithya Ravi, Mehar Bhatia, Maria Antoniak, Yulia Tsvetkov, Vered Shwartz, and Yejin Choi. CulturalBench: A robust, diverse and challenging benchmark for measuring LMs’ cultural knowledge through human-AI red-teaming. In Wanxiang Che, Joyce Nabende, Ekaterina Shutova, and Mohammad Taher Pilehvar (eds.), *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 25663–25701, Vienna, Austria, jul 2025. Association for Computational Linguistics. ISBN 979-8-89176-251-0. doi: 10.18653/v1/2025.acl-long.1247. URL <https://aclanthology.org/2025.acl-long.1247/>.
- Damai Dai, Li Dong, Yaru Hao, Zhifang Sui, Baobao Chang, and Furu Wei. Knowledge neurons in pretrained transformers. In Smaranda Muresan, Preslav Nakov, and Aline Villavicencio (eds.), *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 8493–8502, Dublin, Ireland, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.acl-long.581. URL <https://aclanthology.org/2022.acl-long.581/>.
- Jia Deng, Tianyi Tang, Yanbin Yin, Wenhao yang, Xin Zhao, and Ji-Rong Wen. Neuron based personality trait induction in large language models. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=LYHEY783Np>.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT: Pre-training of deep bidirectional transformers for language understanding. In Jill Burstein, Christy Doran, and Tamar Solorio (eds.), *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pp. 4171–4186, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1423. URL <https://aclanthology.org/N19-1423/>.
- Gemma Team. Gemma 3. 2025. URL <https://goo.gle/Gemma3Report>.
- Mor Geva, Roei Schuster, Jonathan Berant, and Omer Levy. Transformer feed-forward layers are key-value memories. In Marie-Francine Moens, Xuanjing Huang, Lucia Specia, and Scott Wen-tau Yih (eds.), *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pp. 5484–5495, Online and Punta Cana, Dominican Republic, November 2021. Association for Computational Linguistics. doi: 10.18653/v1/2021.emnlp-main.446. URL <https://aclanthology.org/2021.emnlp-main.446/>.
- Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, et al. The llama 3 herd of models, 2024. URL <https://arxiv.org/abs/2407.21783>. version 3.

- C. Haerpfer, R. Inglehart, A. Moreno, C. Welzel, K. Kizilova, J. Diez-Medrano, M. Lagos, P. Norris, E. Ponarin, B. Puranen, et al. World values survey: Round seven – country-pooled datafile, 2020. Madrid, Spain & Vienna, Austria: JD Systems Institute & WWSA Secretariat, doi.org/10.14281/18241.1.
- Takeshi Kojima, Itsuki Okimura, Yusuke Iwasawa, Hitomi Yanaka, and Yutaka Matsuo. On the multilingual ability of decoder-based pre-trained language models: Finding and controlling language-specific neurons. In Kevin Duh, Helena Gomez, and Steven Bethard (eds.), *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 6919–6971, Mexico City, Mexico, June 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.naacl-long.384. URL <https://aclanthology.org/2024.naacl-long.384/>.
- Cheng Li, Mengzhuo Chen, Jindong Wang, Sunayana Sitaram, and Xing Xie. Culturellm: Incorporating cultural differences into large language models. In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang (eds.), *Advances in Neural Information Processing Systems*, volume 37, pp. 84799–84838. Curran Associates, Inc., 2024a. URL [https://proceedings.neurips.cc/paper\\_files/paper/2024/file/9a16935bf54c4af233e25d998b7f4a2c-Paper-Conference.pdf](https://proceedings.neurips.cc/paper_files/paper/2024/file/9a16935bf54c4af233e25d998b7f4a2c-Paper-Conference.pdf).
- Cheng Li, Damien Teney, Linyi Yang, Qingsong Wen, Xing Xie, and Jindong Wang. Culturepark: Boosting cross-cultural understanding in large language models. In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang (eds.), *Advances in Neural Information Processing Systems*, volume 37, pp. 65183–65216. Curran Associates, Inc., 2024b. URL [https://proceedings.neurips.cc/paper\\_files/paper/2024/file/77f089cd16dbc36ddd1cae18446fbdd-Paper-Conference.pdf](https://proceedings.neurips.cc/paper_files/paper/2024/file/77f089cd16dbc36ddd1cae18446fbdd-Paper-Conference.pdf).
- Chen Cecilia Liu, Anna Korhonen, and Iryna Gurevych. Cultural learning-based culture adaptation of language models. In Wanxiang Che, Joyce Nabende, Ekaterina Shutova, and Mohammad Taher Pilehvar (eds.), *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 3114–3134, Vienna, Austria, jul 2025. Association for Computational Linguistics. ISBN 979-8-89176-251-0. doi: 10.18653/v1/2025.acl-long.156. URL <https://aclanthology.org/2025.acl-long.156/>.
- Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=Bkg6RiCqY7>.
- Yun Luo, Zhen Yang, Fandong Meng, Yafu Li, Jie Zhou, and Yue Zhang. An empirical study of catastrophic forgetting in large language models during continual fine-tuning. *arXiv preprint arXiv:2308.08747*, 2025. URL <https://arxiv.org/abs/2308.08747>. version 5.
- Kevin Meng, David Bau, Alex Andonian, and Yonatan Belinkov. Locating and editing factual associations in gpt. In S. Koyejo, S. Mohamed, A. Agarwal, D. Belgrave, K. Cho, and A. Oh (eds.), *Advances in Neural Information Processing Systems*, volume 35, pp. 17359–17372. Curran Associates, Inc., 2022. URL [https://proceedings.neurips.cc/paper\\_files/paper/2022/file/6f1d43d5a82a37e89b0665b33bf3a182-Paper-Conference.pdf](https://proceedings.neurips.cc/paper_files/paper/2022/file/6f1d43d5a82a37e89b0665b33bf3a182-Paper-Conference.pdf).
- Junho Myung, Nayeon Lee, Yi Zhou, Jiho Jin, Rifki Afina Putri, Dimosthenis Antypas, Hsuvas Borkakoty, Eunsu Kim, Carla Perez-Almendros, Abinew Ali Ayele, Victor Gutierrez Basulto, Yazmin Ibanez-Garcia, Hwaran Lee, Shamsuddeen Hassan Muhammad, Kiwoong Park, Anar Sabuhi Rzayev, Nina White, Seid Muhie Yimam, Mohammad Taher Pilehvar, Nedjma Ousidhoum, Jose Camacho-Collados, and Alice Oh. BLEND: A benchmark for LLMs on everyday knowledge in diverse cultures and languages. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024. URL <https://openreview.net/forum?id=nrEqH502eC>.
- Danial Namazifard and Lukas Galke. Isolating culture neurons in multilingual large language models. *arXiv preprint arXiv:2508.02241*, 2025. URL <https://arxiv.org/abs/2508.02241>.

- Tarek Naous, Michael J Ryan, Alan Ritter, and Wei Xu. Having beer after prayer? measuring cultural bias in large language models. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar (eds.), *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 16366–16393, Bangkok, Thailand, aug 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.862. URL <https://aclanthology.org/2024.acl-long.862/>.
- Francesco Ortu, Zhijing Jin, Diego Doimo, Mrinmaya Sachan, Alberto Cazzaniga, and Bernhard Schölkopf. Competition of mechanisms: Tracing how language models handle facts and counterfactuals. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar (eds.), *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 8420–8436, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.458. URL <https://aclanthology.org/2024.acl-long.458/>.
- Qwen Team. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*, 2025. URL <https://arxiv.org/abs/2505.09388>. version 1.
- Abhinav Sukumar Rao, Akhila Yerukola, Vishwa Shah, Katharina Reinecke, and Maarten Sap. NormAd: A framework for measuring the cultural adaptability of large language models. In Luis Chiruzzo, Alan Ritter, and Lu Wang (eds.), *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 2373–2403, Albuquerque, New Mexico, apr 2025. Association for Computational Linguistics. ISBN 979-8-89176-189-6. doi: 10.18653/v1/2025.naacl-long.120. URL <https://aclanthology.org/2025.naacl-long.120/>.
- Noam Shazeer. Glu variants improve transformer. *arXiv preprint arXiv:2002.05202*, 2020. URL <https://arxiv.org/abs/2002.05202>. version 1.
- Nicholas Sukiennik, Chen Gao, Fengli Xu, and Yong Li. An evaluation of cultural value alignment in llm. *arXiv preprint arXiv:2504.08863*, 2025. URL <https://arxiv.org/abs/2504.08863>.
- Alon Talmor, Jonathan Herzig, Nicholas Lourie, and Jonathan Berant. CommonsenseQA: A question answering challenge targeting commonsense knowledge. In Jill Burstein, Christy Doran, and Tamar Solorio (eds.), *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pp. 4149–4158, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1421. URL <https://aclanthology.org/N19-1421/>.
- Tianyi Tang, Wenyang Luo, Haoyang Huang, Dongdong Zhang, Xiaolei Wang, Xin Zhao, Furu Wei, and Ji-Rong Wen. Language-specific neurons: The key to multilingual capabilities in large language models. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar (eds.), *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 5701–5715, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.309. URL <https://aclanthology.org/2024.acl-long.309/>.
- TII Team. The falcon 3 family of open models, December 2024.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL [https://proceedings.neurips.cc/paper\\_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf).
- Alex Wang, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel R. Bowman. GLUE: A multi-task benchmark and analysis platform for natural language understanding. In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=rJ4km2R5t7>.

Wikimedia Foundation. Wikimedia downloads. URL <https://dumps.wikimedia.org>.

Shaoyang Xu, Yongqi Leng, Linhao Yu, and Deyi Xiong. Self-pluralising culture alignment for large language models. In Luis Chiruzzo, Alan Ritter, and Lu Wang (eds.), *Proceedings of the 2025 Conference of the Nations of the Americas Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pp. 6859–6877, Albuquerque, New Mexico, April 2025. Association for Computational Linguistics. ISBN 979-8-89176-189-6. doi: 10.18653/v1/2025.naacl-long.350. URL <https://aclanthology.org/2025.naacl-long.350/>.

Nakyeong Yang, Taegwan Kang, Stanley Jungkyu Choi, Honglak Lee, and Kyomin Jung. Mitigating biases for instruction-following language models via bias neurons elimination. In Lun-Wei Ku, Andre Martins, and Vivek Srikumar (eds.), *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 9061–9073, Bangkok, Thailand, August 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.acl-long.490. URL <https://aclanthology.org/2024.acl-long.490/>.

Jiahao Ying, Wei Tang, Yiran Zhao, Yixin Cao, Yu Rong, and Wenxuan Zhang. Disentangling language and culture for evaluating multilingual large language models. In Wanxiang Che, Joyce Nabende, Ekaterina Shutova, and Mohammad Taher Pilehvar (eds.), *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 22230–22251, Vienna, Austria, July 2025. Association for Computational Linguistics. ISBN 979-8-89176-251-0. doi: 10.18653/v1/2025.acl-long.1082. URL <https://aclanthology.org/2025.acl-long.1082/>.

Zeping Yu and Sophia Ananiadou. Neuron-level knowledge attribution in large language models. In Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen (eds.), *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pp. 3267–3280, Miami, Florida, USA, November 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.emnlp-main.191. URL <https://aclanthology.org/2024.emnlp-main.191/>.

Pengwei Zhan, Zhen Xu, Qian Tan, Jie Song, and Ru Xie. Unveiling the lexical sensitivity of LLMs: Combinatorial optimization for prompt enhancement. In Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen (eds.), *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pp. 5128–5154, Miami, Florida, USA, November 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.emnlp-main.295. URL <https://aclanthology.org/2024.emnlp-main.295/>.

Wenlong Zhao, Debanjan Mondal, Niket Tandon, Danica Dillion, Kurt Gray, and Yuling Gu. World-ValuesBench: A large-scale benchmark dataset for multi-cultural value awareness of language models. In Nicoletta Calzolari, Min-Yen Kan, Veronique Hoste, Alessandro Lenci, Sakriani Sakti, and Nianwen Xue (eds.), *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024)*, pp. 17696–17706, Torino, Italia, may 2024. ELRA and ICCL. URL <https://aclanthology.org/2024.lrec-main.1539/>.

## A DATASET DETAILS

Table 3: Cultural benchmarks used in our experiments.

| Benchmark                                            | #Country | #Instance               | Target                                     |
|------------------------------------------------------|----------|-------------------------|--------------------------------------------|
| BLEnD <sup>3</sup><br>(Myung et al., 2024)           | 16       | 306k MCQs,<br>15k SAQs  | Everyday knowledge in a<br>diverse culture |
| CulturalBench <sup>4</sup><br>(Chiu et al., 2025)    | 45       | 1.23k                   | Cultural knowledge                         |
| NormAd <sup>5</sup><br>(Rao et al., 2025)            | 75       | 2.63k                   | Cultural etiquette and norms               |
| WorldValuesBench <sup>6</sup><br>(Zhao et al., 2024) | ~64      | 260 per<br>participants | Cultural values                            |

Table 3 lists the cultural benchmarks used in our experiments. As explained in Section 2.1, BLEnD evaluates everyday cultural knowledge of LLMs in multiple-choice questions (MCQs) and short answer questions (SAQs). CulturalBench has two task formats: CulturalBench-Easy, which asks about cultural knowledge in multiple-choice questions with four options, and CulturalBench-Hard, which asks whether each of these four options is correct or not with the same question. For simplicity, we adopt CulturalBench-Easy for evaluation. In NormAd, a model is asked to determine whether a given daily scenario is acceptable in a specified culture, which requires understanding of both cultural knowledge and values. The task of WorldValuesBench is to predict participants’ responses to questionnaires given their demographic information. Questionnaires are common for all participants, such as *do you believe in God?*, derived from World Values Survey Wave 7. We download the datasets from Hugging Face Datasets and the GitHub repositories.

Table 4: Examples of BLEnD<sub>neur</sub> and BLEnD<sub>ctrl</sub>.

| BLEnD <sub>neur</sub>                                                                                                                                                                                                                          | BLEnD <sub>ctrl</sub>                                                                                                                                                                     |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| What is a common snack for preschool kids in the UK? Without any explanation, choose only one from the given alphabet choices(e.g., A, B, C). Provide as JSON format:<br>{“answer_choice”:””}<br>A. cookie B. egg C. fruit D. jelly<br>Answer: | Without any explanation, choose only one from the given alphabet choices(e.g., A, B, C). Provide as JSON format:<br>{“answer_choice”:””}<br>A. cookie B. egg C. fruit D. jelly<br>Answer: |

As described in Section 3.3, we prepare BLEnD<sub>ctrl</sub> corresponding to each question of BLEnD<sub>neur</sub>. Table 4 shows examples of BLEnD<sub>neur</sub> and BLEnD<sub>ctrl</sub>. BLEnD<sub>ctrl</sub> is created by omitting the question content from the instances of BLEnD<sub>neur</sub>. In CULNIG, by subtracting the neuron attribution score of BLEnD<sub>ctrl</sub> from that of BLEnD<sub>neur</sub>, we can measure the sheer contribution of neurons to culture knowledge.

Moreover, we constructed the CountryRC (CRC) dataset to filter out superficial neurons that respond to country names. We utilized ChatGPT to create CRC. We instructed ChatGPT to generate reading comprehension problems that contain a country name in their context, and the answer is that country name. We also specified that the problems must not require any cultural understanding. CRC has 50 instances, and 30 instances have only one country name in their context, and the remaining 20 contain an additional dummy country name. Each instance has four answer choices of country names. Country names are represented as placeholders and replaced with the actual names of the target countries.

<sup>3</sup><https://huggingface.co/datasets/nayeon212/BLEnD>

<sup>4</sup><https://huggingface.co/datasets/kellycyy/CulturalBench>

<sup>5</sup><https://huggingface.co/datasets/akhilayerukola/NormAd>

<sup>6</sup><https://github.com/Demon702/WorldValuesBench/tree/635db7455e2c656978929210eba984bc09ddd659>

In addition to the cultural benchmarks, we use NLU benchmarks that do not necessarily require cultural understanding for evaluation. If masking identified neurons results in the reduction of NLU abilities, we cannot conclude that the identified neurons really support the cultural mechanisms of an LLM, even if masking them degrades cultural understanding. We use CommonsenseQA<sup>7</sup>, QNLI<sup>8</sup>, and MRPC<sup>9</sup>. As for CommonsenseQA and QNLI, ground truths for the test set have not been published, so we use the validation set for evaluation.

Considering the sensitivity of LLMs to prompt wording, we prepared four task instructions for each evaluation dataset except for BLENd. For BLENd, task instructions are already included in data sources, and each problem has multiple instances with diverse answer choices, so we used them without additional instructions. We used the prompts in the original paper as a seed and utilized ChatGPT to rephrase the prompts. We show the prompts for each dataset in Table 5, Table 6, Table 7, Table 8, Table 9, Table 10, and Table 11.

Table 5: Evaluation prompts for CulturalBench

| Instruction                                                                                                                                                                                                                    |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| To answer the following multiple-choice question, you should choose one option only among A,B,C,D. Do not output any other things.<br>Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d} |
| Select only one option from A, B, C, or D to answer the following multiple-choice question. Do not output anything else.<br>Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}           |
| Choose one answer among A, B, C, and D for the question below. Do not include any explanation or extra content.<br>Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}                    |
| You must answer the following question by selecting a single choice from A, B, C, or D. Do not write anything else.<br>Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}                |

<sup>7</sup>[https://huggingface.co/datasets/tau/commonsense\\_qa](https://huggingface.co/datasets/tau/commonsense_qa)

<sup>8</sup><https://huggingface.co/datasets/nyu-ml1/glue>

<sup>9</sup><https://huggingface.co/datasets/nyu-ml1/glue>

Table 6: Evaluation prompts for NormAd

| Instruction                                                                                                                                                                                                                                                                                                                                                                                                                                                              |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p>You will be given a short story involving 2 or more characters about everyday situations. You will also be given a country of relevance. Read the story and determine if the character’s behavior aligns with the social norms of that country. Only consider the provided context and country. Do not include any outside assumptions. Respond <b>**only**</b> with the number 1, 2, or 3.</p> <p>country: {country}<br/> Story: {story}<br/> Options: {options}</p> |
| <p>Read a short story describing everyday events with at least two characters. You’ll also receive a country. Your task is to judge whether the character’s actions are in line with the cultural expectations of that country. Use only the given story and cultural context; do not bring in any outside knowledge. Answer strictly with the number 1, 2, or 3.</p> <p>country: {country}<br/> Story: {story}<br/> Options: {options}</p>                              |
| <p>A short scenario featuring multiple characters will be provided, along with the cultural background of a specific country. Your role is to determine how well the behavior shown fits that country’s social norms. Do not make inferences beyond the given content. Only respond with 1, 2, or 3.</p> <p>country: {country}<br/> Story: {story}<br/> Options: {options}</p>                                                                                           |
| <p>Given a short daily-life story involving multiple characters, along with the country, assess whether the behavior depicted fits within the social norms of that culture. Stick strictly to the given material without adding outside reasoning. Answer with just 1, 2, or 3.</p> <p>country: {country}<br/> Story: {story}<br/> Options: {options}</p>                                                                                                                |

Table 7: Evaluation prompts for WorldValuesBench

| Instruction                                                                                                                                                                                                  |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p>System: You are a person from {country}.</p> <p>Prompt: Question: {question}</p> <p>Please respond with a single digit only from {min_option} to {max_option}. Do not include any other text.</p>         |
| <p>System: You are a person from {country}.</p> <p>Prompt: Question: {question}</p> <p>Your answer should be a single digit between {min_option} and {max_option}. Do not add any other information.</p>     |
| <p>System: Behave as if you are from {country}.</p> <p>Prompt: Question: {question}</p> <p>Please respond with a single digit only from {min_option} to {max_option}. Do not include any other text.</p>     |
| <p>System: Behave as if you are from {country}.</p> <p>Prompt: Question: {question}</p> <p>Your answer should be a single digit between {min_option} and {max_option}. Do not add any other information.</p> |

Table 8: Evaluation prompts for CountryRC

| Instruction                                                                                                                                                              |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Read the passage carefully and choose a single option from A, B, C, D to answer the question. Do not output any other text.                                              |
| passage: {passage}<br>question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}                                                           |
| Read the following passage and question. Then, pick the most suitable answer from the four options. Only return the letter of your choice (A, B, C, or D).               |
| passage: {passage}<br>question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}                                                           |
| From the information provided in the passage, choose the best answer to the question. You must select a single choice: 1, 2, 3, or 4, and do not include any other text. |
| passage: {passage}<br>question: {question}<br>1. {option_a}<br>2. {option_b}<br>3. {option_c}<br>4. {option_d}                                                           |
| Determine the correct answer to the question based on the content of the passage. Respond with one of the following: 1, 2, 3, or 4. No additional text is needed.        |
| passage: {passage}<br>question: {question}<br>1. {option_a}<br>2. {option_b}<br>3. {option_c}<br>4. {option_d}                                                           |

Table 9: Evaluation prompts for CommonsenseQA

| Instruction                                                                                                                          |
|--------------------------------------------------------------------------------------------------------------------------------------|
| To answer the following multiple-choice question, you should choose one option only among A,B,C,D,E. Do not output any other things. |
| Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}<br>E. {option_e}                            |
| Choose one answer among A, B, C, D, and E for the question below. Do not include any explanation or extra content.                   |
| Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}<br>E. {option_e}                            |
| Pick one option only — A, B, C, D, or E — as the answer to the question below. Do not provide any additional text.                   |
| Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}<br>E. {option_e}                            |
| Please choose one and only one of the following options (A, B, C, D, or E) to answer the question. Do not add anything else.         |
| Question: {question}<br>A. {option_a}<br>B. {option_b}<br>C. {option_c}<br>D. {option_d}<br>E. {option_e}                            |

Table 10: Evaluation prompts for QNLI

| Instruction                                                                                                                                                                                                                                                                                                                                                                                 |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p>Determine whether the following context sentence contains enough information to answer the question.<br/>           Question: {question}<br/>           Context: {sentence}<br/>           Respond with:<br/>           0 if it does (entailment)<br/>           1 if it does not (not_entailment)<br/>           Only answer with 0 or 1.</p>                                           |
| <p>Classify the relationship between the following question and context.<br/>           Question: {question}<br/>           Context: {sentence}<br/>           Label as:<br/>           0: entailment – the question is supported by the context<br/>           1: not_entailment – the question is not supported by the context<br/>           Please respond with either 0 or 1 only.</p> |
| <p>Read the question and the context.<br/>           Question: {question}<br/>           Context: {sentence}<br/>           If the context provides enough evidence to answer the question, return 0 (entailment).<br/>           If the context is insufficient or irrelevant, return 1 (not_entailment).<br/>           Your answer should be either 0 or 1.</p>                          |
| <p>Your task is to judge if the answer to the question can be found in the context.<br/>           Question: {question}<br/>           Context: {sentence}<br/>           Answer 0 for entailment, and 1 for not_entailment. Do not include any other text.</p>                                                                                                                             |

Table 11: Evaluation prompts for MRPC

| Instruction                                                                                                                                                                                                                                                                                                                                                                                                                              |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <p>Determine whether the following two sentences are paraphrases of each other in meaning.<br/>           Sentence 1: {sentence1}<br/>           Sentence 2: {sentence2}<br/>           Respond with:<br/>           1 – if they are paraphrases<br/>           0 – if they are not paraphrases<br/>           Only answer with 0 or 1.</p>                                                                                              |
| <p>You are given two sentences. Judge whether they express the same meaning, even if the wording is different.<br/>           Sentence 1: {sentence1}<br/>           Sentence 2: {sentence2}<br/>           Answer with 1 if they are paraphrases, and 0 if they are not.<br/>           Please respond using only 0 or 1.</p>                                                                                                           |
| <p>A paraphrase means that two sentences convey the same information using different words or structure.<br/>           Sentence 1: {sentence1}<br/>           Sentence 2: {sentence2}<br/>           Decide whether these sentences are paraphrases.<br/>           Return 1 for paraphrase, 0 for not paraphrase.<br/>           Your answer must be either 0 or 1.</p>                                                                |
| <p>Compare the following two sentences. If they convey the same meaning regardless of differences in wording, classify them as paraphrases.<br/>           Sentence 1: {sentence1}<br/>           Sentence 2: {sentence2}<br/>           Respond with:<br/>           1 – if they are semantically equivalent (paraphrase)<br/>           0 – if they are not semantically equivalent<br/>           Only use 0 or 1 as your answer.</p> |

## B MODEL DETAILS

In our experiment, we analyze six open source state-of-the-art LLMs: gemma-3-12b-it<sup>10</sup>, gemma-3-27b-it<sup>11</sup>, Qwen3-14B<sup>12</sup>, Llama-3-8B-Instruct<sup>13</sup>, phi-4<sup>14</sup>, and Falcon3-10B-Instruct<sup>15</sup>. We apply our methods to these models to show the robustness and generalizability of our findings. We download the parameters from Hugging Face Hub. The architectures of these models are similar, as described in Section 3.1.

Table 12: The total number of neurons in each module of each model.

| Models                | Total Neuron Count |                 |               |                 |
|-----------------------|--------------------|-----------------|---------------|-----------------|
|                       | MLP gate           | Attention query | Attention key | Attention value |
| gemma-3-12b-it        | 737,280            | 196,608         | 98,304        | 98,304          |
| gemma-3-27b-it        | 1,333,248          | 253,952         | 126,976       | 126,976         |
| Qwen3-14B             | 696,320            | 204,800         | 40,960        | 40,690          |
| Llama-3.1-8B-Instruct | 458,752            | 131,072         | 32,768        | 32,768          |
| phi-4                 | 716,800            | 204,800         | 51,200        | 51,200          |
| Falcon3-10B-Instruct  | 921,600            | 122,880         | 40,960        | 40,960          |

The total number of neurons in each model module is shown in Table 12. We only include the modules from which we select culture neurons (see Section 3.1). The number of neurons in an MLP gate module is  $intermediate\_size \times num\_layer$ , the number of neurons in an attention query module is  $head\_dim \times num\_head \times num\_layer$ , and the number of neurons in an attention key and value module is both  $head\_dim \times num\_kv\_head \times num\_layer$ . When using grouped-query attention of the group size  $g$ ,  $num\_kv\_head = num\_head \div g$ .

<sup>10</sup><https://huggingface.co/google/gemma-3-12b-it>

<sup>11</sup><https://huggingface.co/google/gemma-3-27b-it>

<sup>12</sup><https://huggingface.co/Qwen/Qwen3-14B>

<sup>13</sup><https://huggingface.co/meta-llama/Llama-3.1-8B-Instruct>

<sup>14</sup><https://huggingface.co/microsoft/phi-4>

<sup>15</sup><https://huggingface.co/tiiuae/Falcon3-10B-Instruct>

## C CULTURE NEURON DISTRIBUTION

We show the distributions of *culture-general* neurons in each model in Figure 3, Figure 6, Figure 7, Figure 8, Figure 9, and Figure 10. We can observe that the neuron distributions are similar for all the models, concentrated in shallow to middle MLP layers. This result suggests that our method captures mechanisms shared across LLMs.

In addition, Figure 11a shows the distribution of Chinese *culture-specific* neurons in gemma-3-12b-it, and Figure 11b shows the distribution of Chinese neurons identified by CAPE (pure). While CULNIG-specific Chinese neurons are mainly located in shallow to middle MLP layers, similarly to CULNIG-general, CAPE Chinese neurons are concentrated in deeper layers.

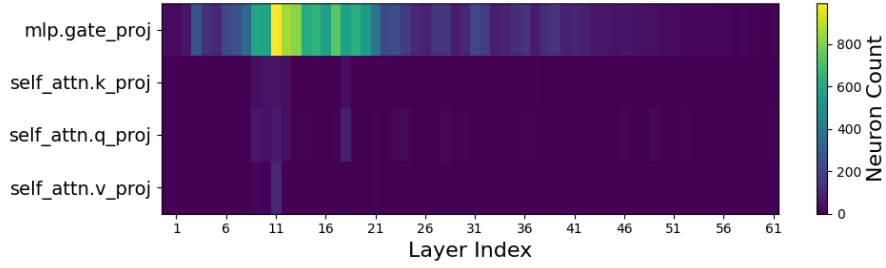


Figure 6: The distribution of *culture-general* neurons in gemma-3-27b-it.

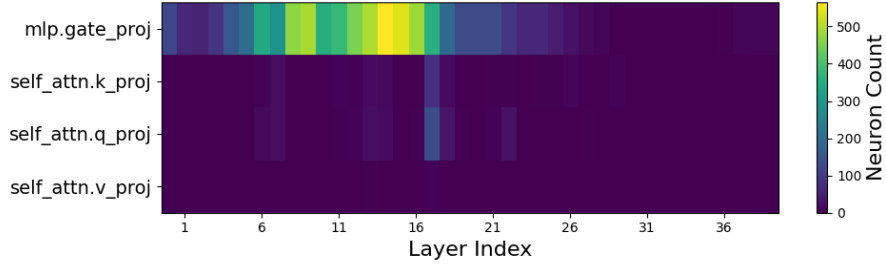


Figure 7: The distribution of *culture-general* neurons in Qwen3-14B.

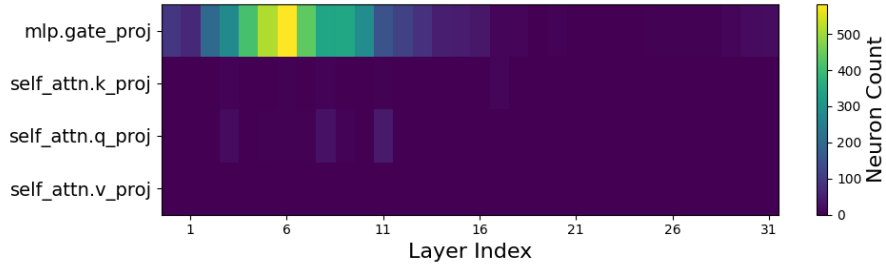
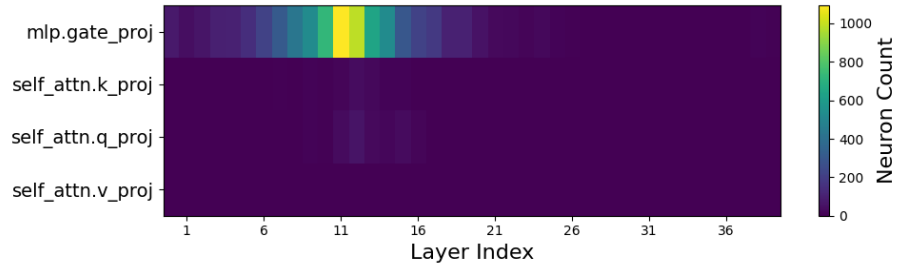
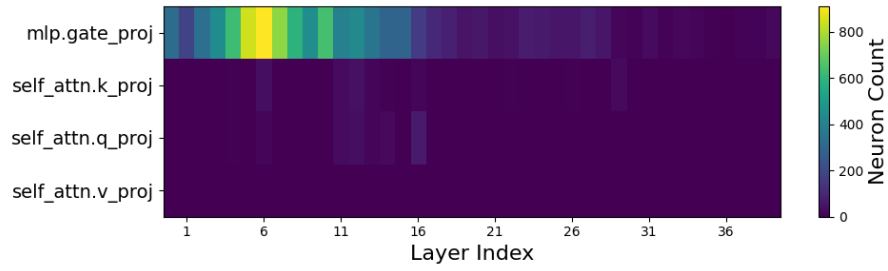
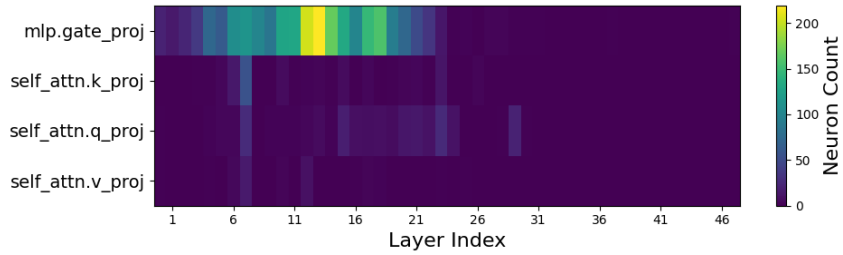
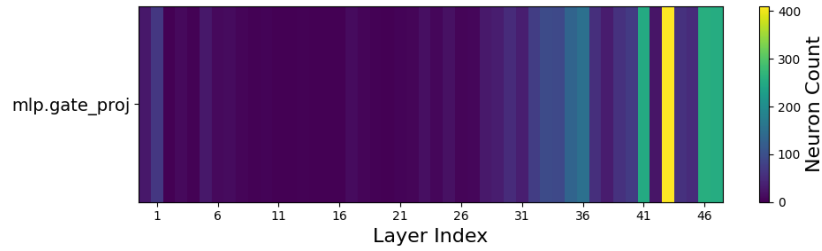


Figure 8: The distribution of *culture-general* neurons in Llama-3.1-8B-Instruct.

Figure 9: The distribution of *culture-general* neurons in phi-4.Figure 10: The distribution of *culture-general* neurons in Falcon3-10B-Instruct.

(a) CULNIG-specific



(b) CAPE

Figure 11: The distribution of Chinese neurons identified by CULNIG-specific and CAPE in gemma-3-12b-it.

## D RESULTS OF MULTILINGUAL EVALUATION ON BLEND SAQ

As explained in Section 2.1, BLEnD provides two types of tasks: multilingual short answer questions (SAQs) and English multiple-choice questions (MCQs). The evaluation results on MCQs are shown in Section 4.3, confirming that suppressing *culture-general* neurons substantially degrades the performance of the models on MCQs (BLEnD<sub>test</sub>). Here, we evaluate LLMs on BLEnD SAQs to see whether *culture-general* neurons are responsible for cultural understanding in multilingual and SAQ settings.

BLEnD covers 16 cultures, and the SAQs for each culture are provided in English and their corresponding language, resulting in 13 languages in total. We prompt LLMs only in their native language to evaluate each culture. Also, to align with the evaluation on MCQs, we use the same three categories as BLEnD<sub>test</sub>. As for the task instruction, we utilize the prompts provided in their GitHub repository<sup>16</sup> and randomly select one instruction per instance. For other details of the evaluation, we follow the original settings of BLEnD (Myung et al., 2024) and their GitHub repositories. We set `max_new_tokens` to 512 and other parameters to the models’ default values. When judging models’ responses, we first lemmatize, stem, or tokenize the models’ responses and the annotation answers. We regard the prediction as correct if any answers are included in the response.

Table 13: Evaluation accuracy (%) on BLEnD SAQs for the original model (Orig), when *culture-general* neurons are masked (Cult), and when random neurons are masked (Rand).

| Model                 | Orig  | Cult  | Rand  |
|-----------------------|-------|-------|-------|
| gemma-3-12b-it        | 51.13 | 42.77 | 49.13 |
| gemma-3-27b-it        | 57.71 | 47.00 | 56.32 |
| Qwen3-14B             | 47.74 | 36.04 | 46.32 |
| Llama-3.1-8B-Instruct | 43.89 | 20.63 | 39.38 |
| phi-4                 | 47.97 | 35.98 | 47.76 |
| Falcon3-10B-Instruct  | 28.36 | 23.41 | 26.91 |

The accuracies on the SAQs are shown in Table 13. Suppressing *culture-general* neurons reduces the accuracy of all models more significantly than suppressing random neurons. Moreover, Figure 12, Figure 13, Figure 14, Figure 15, Figure 16, and Figure 17 show culture-wise accuracies of each model. We can observe that score reduction occurs regardless of cultures. These results indicate that *culture-general* neurons contribute to cultural understanding in multilingual and SAQ settings as well.

<sup>16</sup><https://github.com/nlee0212/BLEnD/tree/9972379c4fd20601691c45e6d7befa6a3eed7ed4>

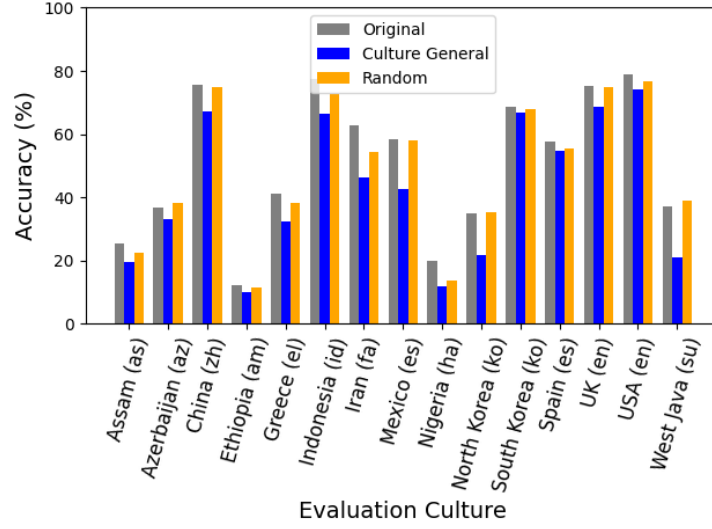


Figure 12: Accuracy of gemma-3-12b-it on BLEnD SAQs for each culture. Evaluation results of the original model, when masking *culture-general* neurons, and when masking random neurons are shown.

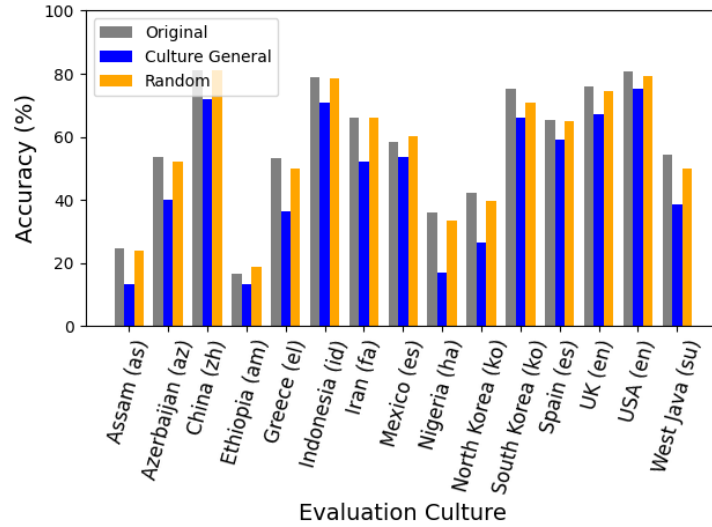


Figure 13: Accuracy of gemma-3-27b-it on BLEnD SAQs for each culture. Evaluation results of the original model, when masking *culture-general* neurons, and when masking random neurons are shown.

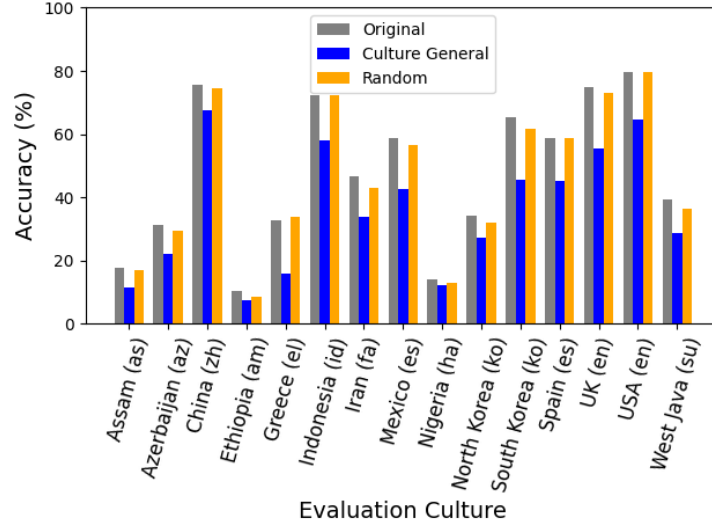


Figure 14: Accuracy of Qwen3-14B on BLEnD SAQs for each culture. Evaluation results of the original model, when masking *culture-general* neurons, and when masking random neurons are shown.

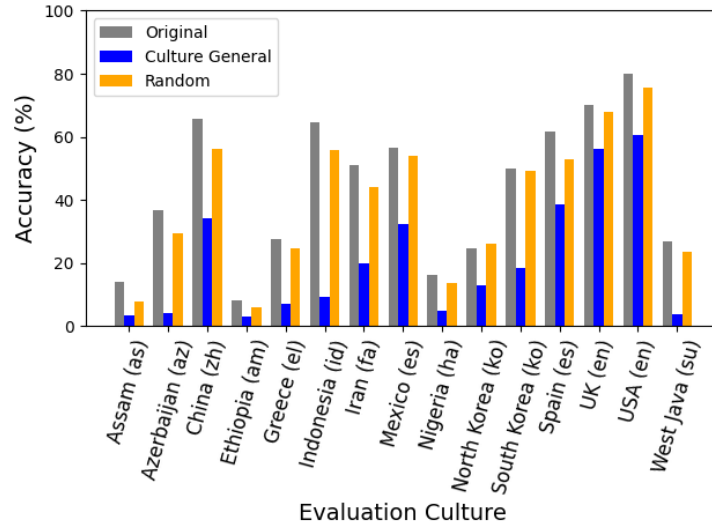


Figure 15: Accuracy of Llama-3.1-8B-Instruct on BLEnD SAQs for each culture. Evaluation results of the original model, when masking *culture-general* neurons, and when masking random neurons are shown.

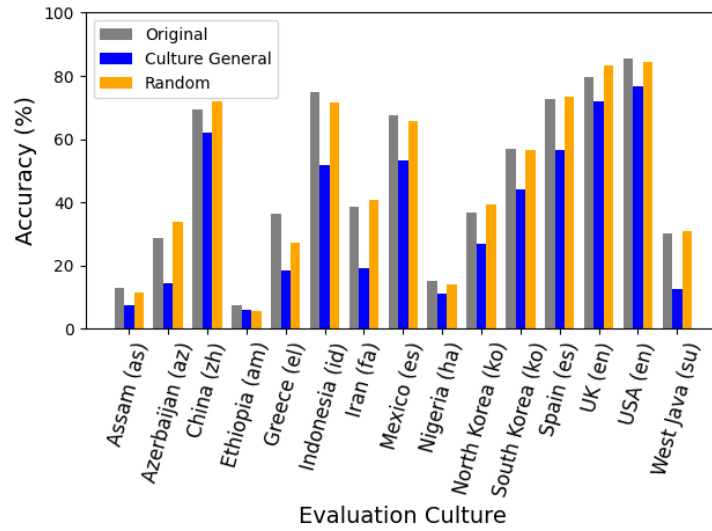


Figure 16: Accuracy of phi-4 on BLEND SAQs for each culture. Evaluation results of the original model, when masking *culture-general* neurons, and when masking random neurons are shown.

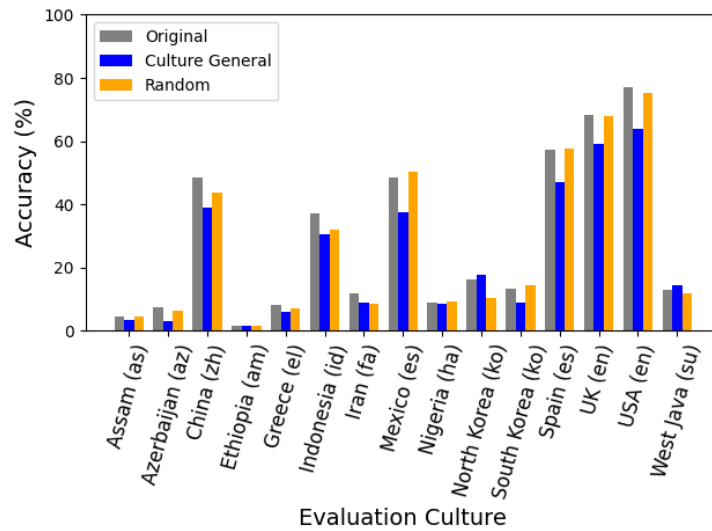


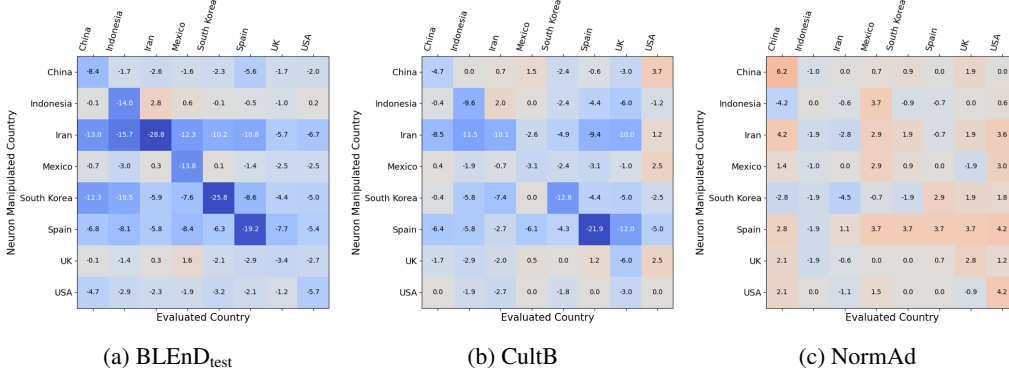
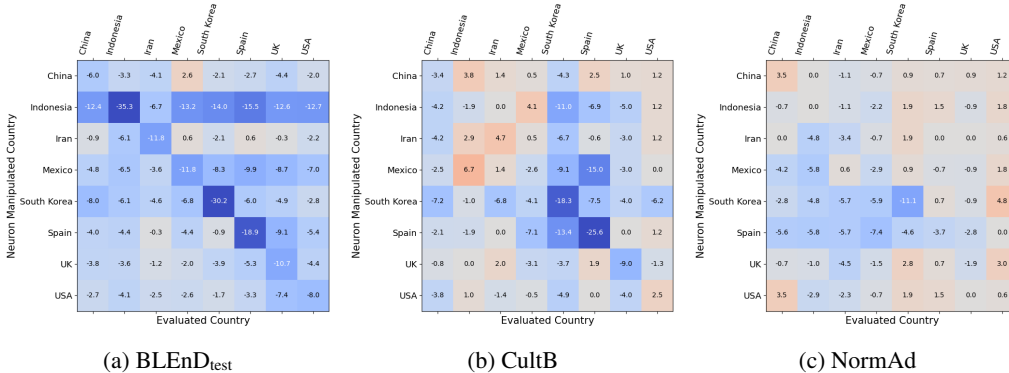
Figure 17: Accuracy of Falcon3-10B-Instruct on BLEND SAQs for each culture. Evaluation results of the original model, when masking *culture-general* neurons, and when masking random neurons are shown.

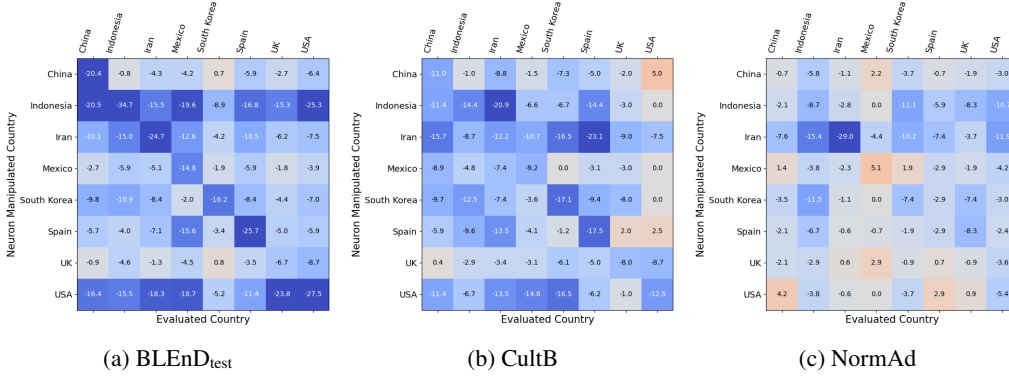
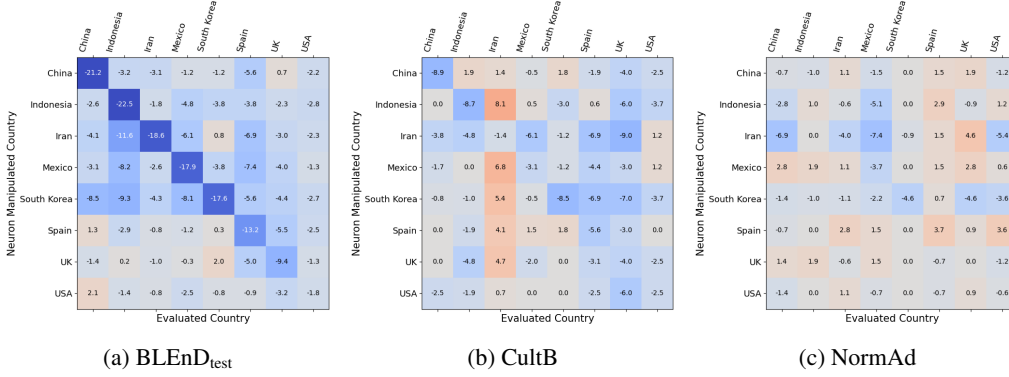
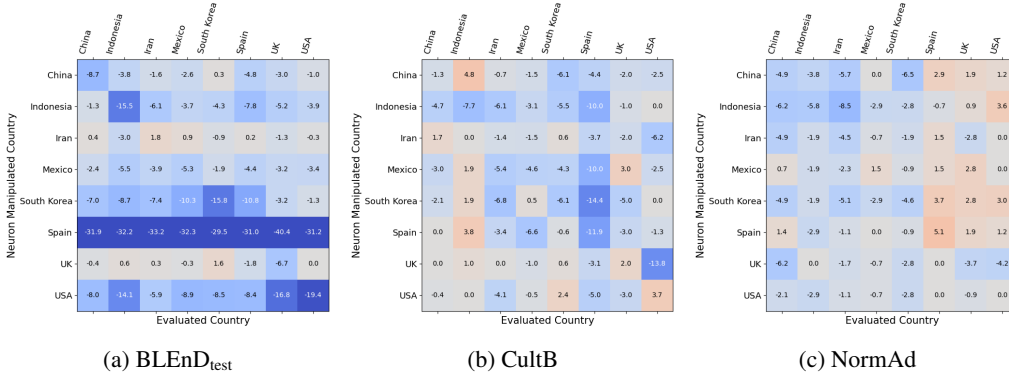
## E RESULTS OF CULTURE-SPECIFIC NEURONS

Table 14: The number of *culture-specific* neurons identified by CULNIG-specific.

| Model                 | China | Indonesia | Iran  | Mexico | South Korea | Spain | UK    | USA   |
|-----------------------|-------|-----------|-------|--------|-------------|-------|-------|-------|
| gemma-3-12b-it        | 2,667 | 2,569     | 2,948 | 2,756  | 2,101       | 2,655 | 2,977 | 3,041 |
| gemma-3-27b-it        | 4,011 | 4,563     | 4,953 | 3,580  | 5,061       | 4,663 | 3,768 | 3,821 |
| Qwen3-14B             | 1,553 | 1,897     | 1,473 | 2,070  | 2,204       | 1,874 | 2,029 | 2,190 |
| Llama-3.1-8B-Instruct | 471   | 782       | 549   | 373    | 678         | 540   | 345   | 470   |
| phi-4                 | 1,072 | 1,192     | 1,373 | 1,249  | 1,524       | 1,039 | 1,210 | 1,050 |
| Falcon3-10B-Instruct  | 1,789 | 2,199     | 1,785 | 1,923  | 2,603       | 2,114 | 1,665 | 2,356 |

In this section, we present the results of *culture-specific* neurons for models not shown in Section 4.4. First, Table 14 shows the number of neurons identified by CULNIG-specific for each country. It is natural that the numbers are proportional to the total number of neurons (see Table 12) because the initial candidate neurons are the top 0.3% neurons ranked by attribution score. Subsequently, *culture-specific* neurons are refined by  $\text{CRC}_{\text{neur}}$  and z-score, which may make the difference between countries. In the table, the number of neurons corresponding to South Korea tends to be large, indicating that models possess more dedicated neurons for South Korean culture than others.

Figure 18: Score reductions after masking *culture-specific* neurons of gemma-3-27b-it.Figure 19: Score reductions after masking *culture-specific* neurons of Qwen3-14B.

Figure 20: Score reductions after masking *culture-specific* neurons of Llama-3.1-8B-Instruct.Figure 21: Score reductions after masking *culture-specific* neurons of phi-4.Figure 22: Score reductions after masking *culture-specific* neurons of Falcon3-10B-Instruct.

We show the evaluation results of *culture-specific* neurons for each model in Figure 4, Figure 18, Figure 19, Figure 20, Figure 21, and Figure 22. These figures show the reduction of scores compared to the original models for each problem culture. The result patterns are similar to Section 4.4. The scores of the same countries as the neuron targets are most affected for BLenDtest and CultB, while a clear pattern is not observed for NormAd. On the other hand, there are several cases where suppressing identified *culture-specific* neurons consistently degrades the scores for all evaluation cultures (e.g., Indonesian neurons in BLenDtest). For these cases, the identified neurons may actually be important for understanding the benchmark task.

Table 15: Average ranks of performance drops among 16 cultures of  $\text{BLEnD}_{\text{test}}$  when masking *culture-specific* neurons. Ranks are averaged over six models.

| Evaluation culture | Neuron culture |             |             |             |             |             |             |             |
|--------------------|----------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                    | China          | Indonesia   | Iran        | Mexico      | South Korea | Spain       | UK          | USA         |
| Algeria            | 14.17          | 9.67        | 8.00        | 10.33       | 12.33       | 13.17       | 7.33        | 8.00        |
| Assam              | 12.17          | 8.50        | 9.83        | 10.83       | 11.00       | 11.17       | 9.83        | 12.00       |
| Azerbaijan         | 5.67           | 11.33       | 4.83        | 8.17        | 9.67        | 8.33        | 7.67        | 8.17        |
| China              | <b>1.00</b>    | 9.83        | 8.33        | 10.00       | 4.33        | 7.00        | 9.83        | 7.33        |
| Ethiopia           | 13.50          | 12.50       | 12.83       | 11.50       | 12.33       | 14.00       | 14.33       | 12.50       |
| Greece             | 8.33           | 10.67       | 8.50        | 12.50       | 8.33        | 8.17        | 7.67        | 8.17        |
| Indonesia          | 8.83           | <b>1.17</b> | 4.00        | 4.00        | 5.33        | 4.67        | 8.00        | 6.33        |
| Iran               | 6.33           | 11.50       | <b>3.50</b> | 10.00       | 9.50        | 8.50        | 11.67       | 8.83        |
| Mexico             | 10.50          | 9.00        | 9.17        | <b>1.17</b> | 6.17        | 4.33        | 10.17       | 6.33        |
| Nigeria            | 7.50           | 7.00        | 6.00        | 9.00        | 11.67       | 12.83       | 9.67        | 11.00       |
| North Korea        | 7.50           | 11.33       | 13.00       | 11.67       | 6.50        | 12.83       | 13.00       | 12.50       |
| South Korea        | 10.83          | 8.33        | 11.33       | 9.50        | <b>1.00</b> | 10.33       | 10.33       | 9.67        |
| Spain              | 3.67           | 7.17        | 9.33        | 3.83        | 5.50        | <b>2.00</b> | 3.17        | 7.83        |
| UK                 | 7.83           | 8.17        | 10.83       | 8.67        | 9.67        | 3.17        | <b>1.17</b> | 5.67        |
| USA                | 8.17           | 7.83        | 11.17       | 7.00        | 11.67       | 7.33        | 5.17        | <b>1.67</b> |
| West Java          | 10.00          | 2.00        | 5.33        | 7.83        | 11.00       | 8.17        | 7.00        | 10.00       |

For a deeper analysis, we show the average ranks of performance drops among 16 cultures of  $\text{BLEnD}_{\text{test}}$  when masking *culture-specific* neurons in Table 15. The ranks are averaged over six models, and a higher rank means a significant drop. We observe that the top ranks are always when the neuron target culture and evaluation culture agree, validating that the identified neurons especially contribute to their target culture. Additionally, when *culture-specific* neurons of a specific culture are masked, it tends to have an impact on scores of related cultures. For example, when Mexican neurons are masked, Spain is the second most strongly influenced culture. When Spanish neurons are masked, Mexico is the third most influenced, and the second most influenced culture is the UK. Spain and Mexico are historically connected, and Spain and the UK are geographically close.

## F REPLICATION OF LAPE AND CAPE

As discussed in Section 4.3 and Appendix C, the neuron distributions of CULNIG and CAPE are different. CAPE is a method designed to isolate culture neurons from language neurons of LAPE, using a multilingual and multicultural dataset, MUREL. We replicate the experiments of LAPE and CAPE according to their GitHub repository<sup>17</sup> and compare with our results.

LAPE identifies language-specific neurons using multilingual corpora taken from Wikipedia<sup>18</sup> (Wikimedia Foundation). Similarly, CAPE first selects neurons using MUREL (in this paper, we call this neuron set “MUREL neuron” to avoid conflict with the name “culture neuron” with CULNIG), and then refines neurons by excluding corresponding LAPE neurons to obtain “pure” culture neurons. As MUREL contains six languages and cultures (Danish (da), German (de), English (en), Persian (fa), Russian (ru), and Chinese (zh)), we identify neurons of these languages and cultures. For the model, we use gemma-3-12b-it.

Table 16: Neuron counts of LAPE and CAPE neurons in gemma-3-12b-it. Languages (cultures) are Danish (da), German (de), English (en), Persian (fa), Russian (ru), and Chinese (zh).

|       | da  | de    | en    | fa    | ru    | zh    |
|-------|-----|-------|-------|-------|-------|-------|
| LAPE  | 914 | 1,087 | 773   | 1,115 | 1,157 | 2,440 |
| MUREL | 412 | 477   | 1,059 | 718   | 810   | 3,897 |
| pure  | 60  | 80    | 644   | 264   | 221   | 2,462 |

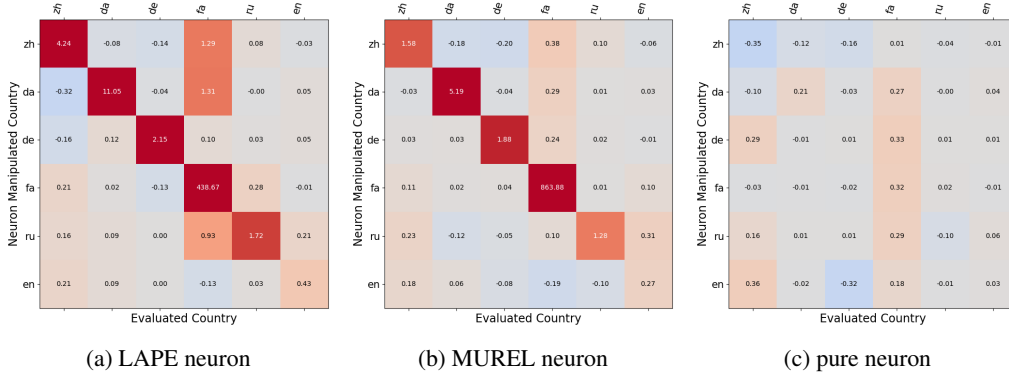


Figure 23: Perplexity increase when masking LAPE, MUREL, and pure neurons from the original state of gemma-3-12b-it.

The number of identified neurons by LAPE and CAPE is shown in Table 16. The number of pure neurons is small, especially for da (60) and de (80). This indicates that the overlaps between LAPE and MUREL neurons are large, failing to isolate culture neurons from language neurons. For evaluation in the CAPE paper, they use the MUREL test set and see the perplexity change. We present the replicated evaluation results in Figure 23. It shows that for LAPE and MUREL neurons, increases in perplexity are most significant when the language or culture of neurons and data match. However, this is not the case for pure neurons, which have little impact after masking them. Based on these results, we speculate that most of the MUREL neurons are actually language neurons. Note that they use gemma-3-12b-pt in the original experiment, while we use gemma-3-12b-it, which is developed by performing instruction tuning on gemma-3-12b-pt, for consistency with our experiment. Other possible differences are hyperparameters, such as the context length of inputs.

Moreover, the evaluation results on BLEnD<sub>test</sub>, CultB, and NormAd for LAPE, MUREL, and pure neurons are presented in Figure 24, Figure 25, and Figure 26, respectively. We show the score

<sup>17</sup>[https://github.com/namazifard/Culture\\_Neurons/tree/f48acc08d2d4a9117610f3e8e29a502fca2704c4](https://github.com/namazifard/Culture_Neurons/tree/f48acc08d2d4a9117610f3e8e29a502fca2704c4)

<sup>18</sup><https://huggingface.co/datasets/wikimedia/wikipedia>

changes from the original model on problems of the cultures common in MUREL and each benchmark. As a result, none of the three methods caused significant changes to the scores. One plausible reason is that all the problems in these benchmarks are asked in English in our evaluation. As shown in Figure 23, the impacts on English are the smallest for all methods. Therefore, if identified neurons contribute to language abilities, not cultural understandings, the effects will be small when asking cultural questions in English.

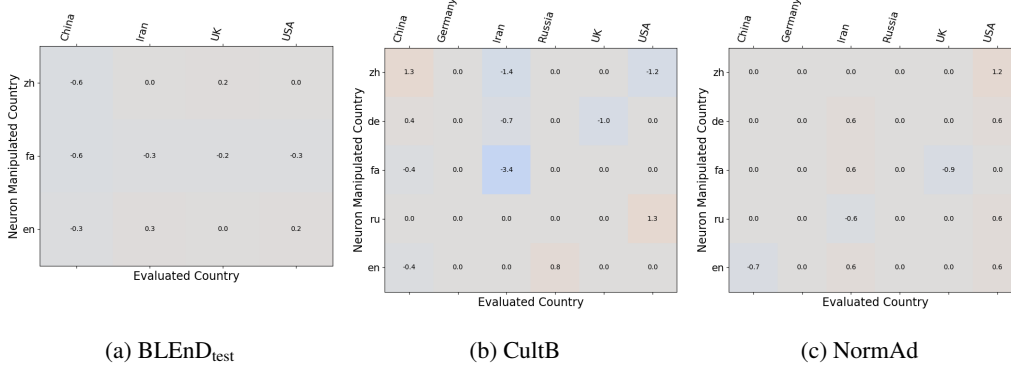


Figure 24: Score reductions after masking LAPE neurons of gemma-3-12b-it.

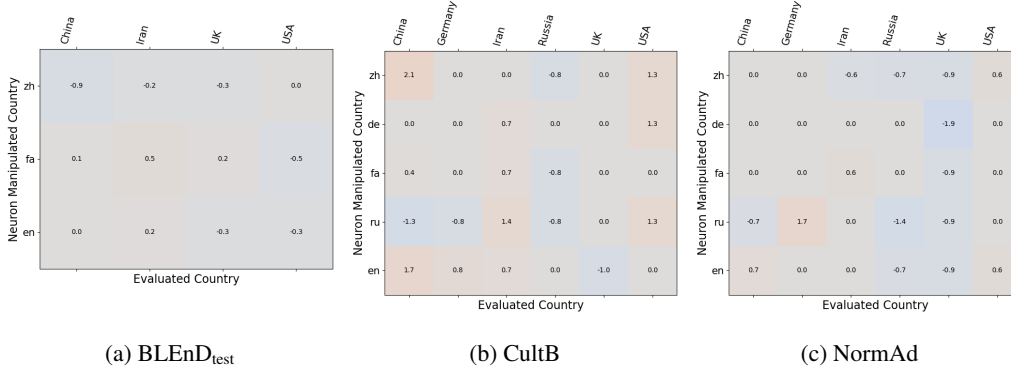


Figure 25: Score reductions after masking MUREL neurons of gemma-3-12b-it.



Figure 26: Score reductions after masking pure neurons of gemma-3-12b-it.

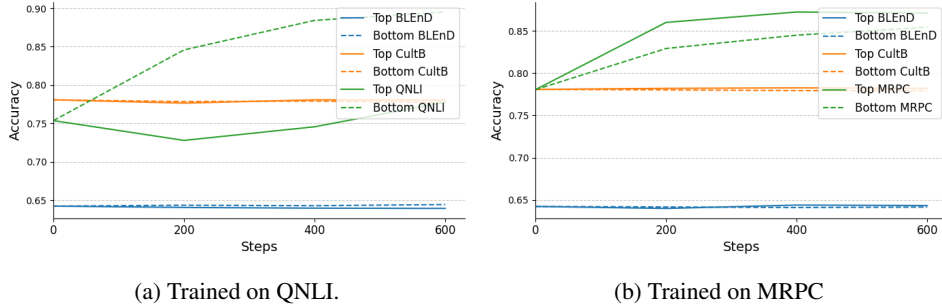
## G DETAILS OF MODEL TRAINING

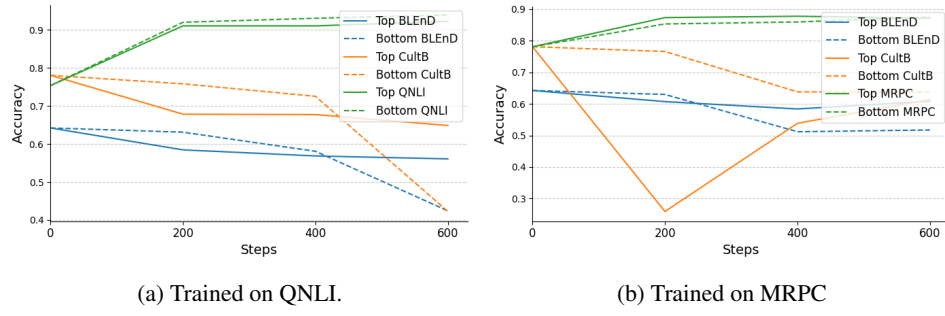
Table 17: Selection of top-culture and bottom-culture modules. Values in parentheses denote the number of culture neurons contained in each module.

| top-culture       | bottom-culture        |
|-------------------|-----------------------|
| layer13 MLP (785) | layer0 attention (0)  |
| layer12 MLP (685) | layer2 attention (0)  |
| layer14 MLP (612) | layer8 attention (0)  |
| layer10 MLP (491) | layer27 attention (0) |
| layer11 MLP (469) | layer28 attention (0) |
| layer7 MLP (429)  | layer30 attention (0) |
| layer9 MLP (409)  | layer31 attention (0) |
|                   | layer32 attention (0) |
|                   | layer36 attention (0) |
|                   | layer39 attention (0) |
|                   | layer41 attention (0) |
|                   | layer43 attention (0) |
|                   | layer44 attention (0) |
|                   | layer45 attention (0) |
|                   | layer47 attention (0) |
|                   | layer34 MLP (0)       |
|                   | layer41 MLP (0)       |
|                   | layer43 MLP (0)       |

In this section, we present the supplementary information and results of Section 4.5. As described in Section 4.5, we select top-culture and bottom-culture modules to fine-tune gemma-3-12b-it with QNLI and MRPC. The selected modules are shown in Table 17. The top-culture modules are all MLP modules of shallow to middle layers, while the bottom-culture modules consist of the shallowest attention, deep attention, and deep MLP modules. This selection matches the distribution of culture-general neurons (subsection 4.3 and Appendix C). Note that for the bottom-culture modules, we randomly picked modules from those without any culture-general neurons to account for 10% of the total parameters. During training, we use AdamW (Loshchilov & Hutter, 2019) optimizer and linear scheduler with batch size 16. For QNLI, we randomly selected 10,000 training samples to reduce computational cost.

The evaluation results when the learning rate is  $3e-5$  are shown in Section 4.5 (Figure 5). We also show the results when the learning rate is  $1e-5$  and  $5e-5$  in Figure 27 and Figure 28, respectively. When the learning rate is  $1e-5$ , there are almost no impacts on cultural benchmarks regardless of updated modules. When the learning rate is  $5e-5$ , the scores on cultural benchmarks degrade at early steps when updating top-culture modules, and the scores also decrease for bottom-culture modules as the training goes on. Regarding the target benchmarks (QNLI or MRPC), the scores improve on all cases except for QNLI when targeting top-culture modules. These results suggest that forgetting can be avoided depending on the learning rate (or other parameters), but tuning bottom-culture modules can achieve a better outcome.

Figure 27: Evaluation results when  $lr=1e-5$ .

Figure 28: Evaluation results when  $lr=5e-5$ .

## H ABLATION OF DATASETS USED FOR NEURON IDENTIFICATION

Table 18: Evaluation results on  $CRC_{test}$  when masking neurons identified with and without  $CRC_{neur}$ .

| Model                 | orig   | w/ $CRC_{neur}$ | w/o $CRC_{neur}$ |
|-----------------------|--------|-----------------|------------------|
| gemma-3-12b-it        | 100.00 | 100.00          | 99.75            |
| gemma-3-27b-it        | 100.00 | 100.00          | 100.00           |
| Qwen3-14B             | 100.00 | 100.00          | 100.00           |
| Llama-3.1-8B-Instruct | 100.00 | 96.00           | 0.00             |
| phi-4                 | 100.00 | 100.00          | 0.13             |
| Falcon3-10B-Instruct  | 100.00 | 99.62           | 98.25            |

CULNIG identifies culture neurons using  $BLEnD_{neur}$ ,  $BLEnD_{ctrl}$ , and  $CRC_{neur}$  (Section 3.3). In this section, we perform the ablation studies of these datasets. Table 18 compares the evaluation results on  $CRC_{test}$  when masking neurons identified by CULNIG-general with and without  $CRC_{test}$ . We observe that without  $CRC_{neur}$ , masking identified neurons significantly reduces accuracy for some models. As described in Section 3.3, we use  $CRC_{neur}$  in CULNIG to eliminate superficial neurons activated by tokens of country names, since such neurons should not be considered as supporting cultural mechanisms. The results confirm that  $CRC_{neur}$  filters out such neurons.

Table 19: Evaluation results of gemma-3-12b-it when masking neurons identified by CULNIG-general with and without  $BLEnD_{ctrl}$ .

|          | #Neuron | $BLEnD_{test}$ | CultB | NormAd | WVB   | ComQA | QNLI  | MRPC  |
|----------|---------|----------------|-------|--------|-------|-------|-------|-------|
| orig     | 0       | 64.22          | 78.08 | 58.54  | 64.08 | 79.71 | 75.37 | 78.04 |
| w/ ctrl  | 8,087   | 37.93          | 62.00 | 52.02  | 58.46 | 75.10 | 72.77 | 78.65 |
| w/o ctrl | 6,494   | 39.65          | 61.57 | 52.82  | 62.28 | 70.13 | 67.49 | 78.25 |

Moreover, Table 19 shows the ablation results of  $BLEnD_{ctrl}$  on gemma-3-12b-it. We observe that without  $BLEnD_{ctrl}$ , the evaluation scores on the NLU benchmarks are worse than normal CULNIG-general, although the number of neurons is smaller. This result indicates that neurons that contribute to properties other than cultural understanding, such as language understanding, tend to get high scores and be selected without  $BLEnD_{ctrl}$ . These results confirm that the datasets used in our pipeline are important for accurately and steadily identifying culture neurons.

## I COMPARISON OF NEURON ATTRIBUTION SCORES

As explained in Section 3.2, we adopt a gradient-based score to measure neuron attribution in solving cultural problems, following Yang et al. (2024). In this section, we compare it with alternative attribution methods.

In our method, the attribution score of a neuron at the  $i$ -th token position is calculated as Equation 6, and aggregated across tokens by taking the maximum (Equation 7). As alternatives, we consider:

- Mean aggregation (*mean*): replacing the maximum with the mean across token positions.
- Weight-gradient inner product (*norm*): directly computing inner product  $\mathbf{w} \cdot \frac{\partial P(y|x)}{\partial \mathbf{w}}$  for the subkey  $\mathbf{w}$  (row vectors of MLP gate, attention query, key, and value modules) associated with each neuron.

Table 20: Evaluation results of masking *culture-general* neurons identified with *max* (the one used in the original pipeline), *mean*, and *norm* attribution scores on gemma-3-12b-it.

| Score       | #Neuron | BLEND <sub>test</sub> | CultB | NormAd | WVB   | ComQA | QNLI  | MRPC  |
|-------------|---------|-----------------------|-------|--------|-------|-------|-------|-------|
| (orig)      | 0       | 64.22                 | 78.08 | 58.54  | 64.08 | 79.71 | 75.37 | 78.04 |
| <i>max</i>  | 8,087   | 37.93                 | 62.00 | 52.02  | 58.46 | 75.10 | 72.77 | 78.65 |
| <i>mean</i> | 8,151   | 59.50                 | 75.75 | 58.59  | 64.27 | 79.20 | 74.22 | 78.23 |
| <i>norm</i> | 8,151   | 56.54                 | 75.06 | 58.86  | 63/72 | 79.71 | 74.86 | 78/20 |

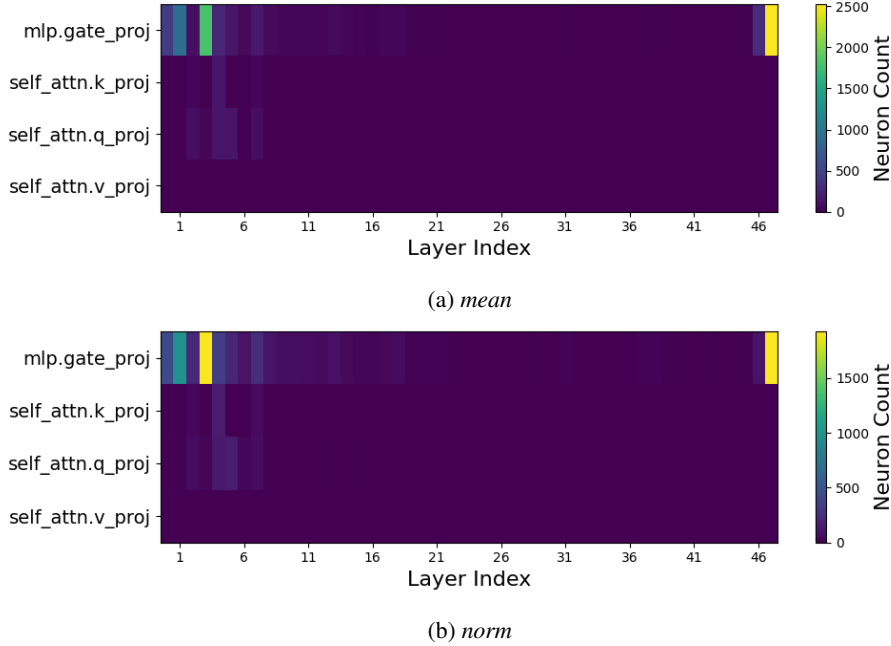


Figure 29: The distribution of neurons identified with *mean* and *norm* attribution scores.

We identify *culture-general* neurons in gemma-3-12b-it with *mean* and *norm* scores integrated into CULNIG-general. The evaluation results are shown in Table 20. Masking *mean* or *norm* barely affects the benchmark scores, indicating that identified neurons do not engage in model behavior. Figure 29 shows that such neurons are mainly located in very shallow and very deep MLP layers, unlike the distribution of our original method (Figure 3). Moreover, Figure 30 compares the distribution of attribution scores. For *max*, the distribution has a wider positive tail, while for *mean* and *norm*, only a few neurons have a positive score. Actually, the number of neurons with z-score  $\geq 2.5$  is 729 for *max*, but only 6 for *mean* and 15 for *norm*, suggesting that *mean* and *norm* failed to

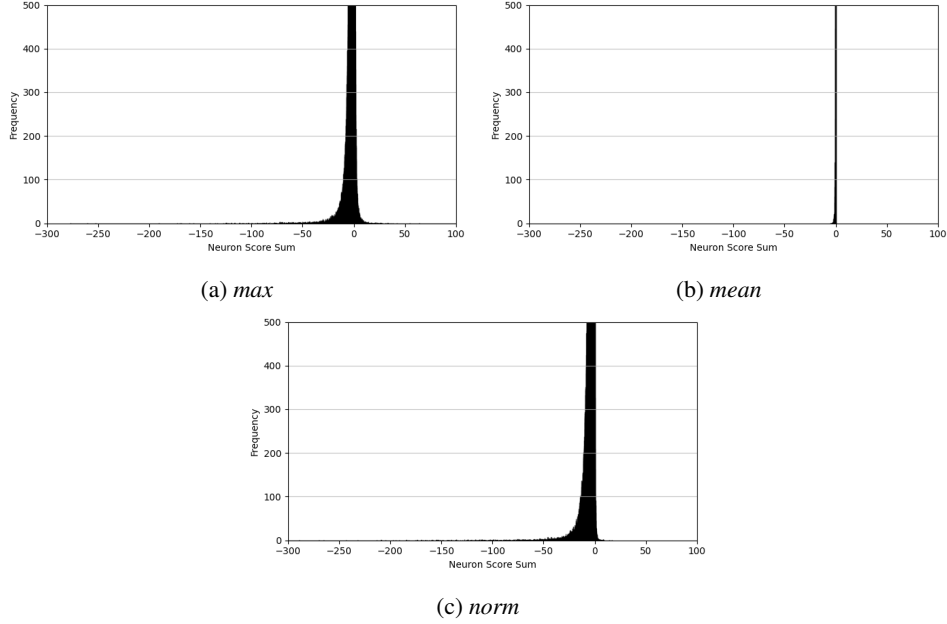


Figure 30: The distribution of neuron attribution scores with *max*, *mean*, and *norm* on gemma-3-12b-it.

distinguish neurons that contribute to cultural understanding. We speculate that this is because the scores of *mean* and *norm* take into account all token positions. Not all tokens necessarily encode cultural representations, so attribution can be obscure. In contrast, *max* highlights salient tokens, which may result in the best performance for identifying culture neurons.

## J EXPERIMENTAL CONFIGURATION

In our experiments, we used NVIDIA H100 GPUs. To calculate all neuron attribution scores on  $\text{BLEnD}_{\text{neur}}$ ,  $\text{BLEnD}_{\text{ctrl}}$  and  $\text{CRC}_{\text{neur}}$  in CULNIG, it took up to 4 hours per model with one H100 GPU (for gemma-3-27b-it, we used two H100 GPUs). For fine-tuning in Section 4.5, it took up to 20 minutes to train gemma-3-12b-it for 600 steps.

## K LLM USAGE

We utilized ChatGPT to construct the CRC dataset and to generate task instructions in the evaluation prompts. We also used ChatGPT and Gemini<sup>19</sup> to proofread the paper. When implementing the scripts for our experiments, we used GitHub Copilot<sup>20</sup> as a coding assistant.

<sup>19</sup><https://gemini.google/about/>

<sup>20</sup><https://github.com/features/copilot>