

Table 1: Additional experimentst on GPT-J and GPT2-XL to illustrate our method’s generalization ability. The mean and standard deviation are reported for 3 repetitions with different ICL examples.

Method	GPT-J					GPT2-XL				
	ICL	CC	DC	PC	UniBias	ICL	CC	DC	PC	UniBias
SST-2	90.67 _{1.06}	90.75 _{1.88}	91.90 _{1.12}	89.14 _{1.68}	92.01 _{1.04}	56.69 _{3.39}	81.35 _{4.36}	87.00 _{2.66}	88.69 _{1.03}	88.69 _{1.10}
WiC	50.21 _{0.99}	51.83 _{0.53}	50.89 _{0.90}	51.67 _{0.59}	52.77 _{0.52}	49.06 _{0.25}	50.88 _{0.92}	50.31 _{1.33}	48.69 _{1.79}	51.22 _{0.69}
COPA	47.33 _{0.47}	47.00 _{0.35}	48.66 _{0.94}	52.99 _{5.88}	55.33 _{2.62}	49.00 _{1.63}	46.33 _{1.70}	49.33 _{1.25}	52.33 _{2.36}	52.00 _{3.74}
MR	89.26 _{1.65}	85.31 _{2.36}	90.29 _{0.45}	89.66 _{0.27}	90.39 _{0.36}	51.11 _{1.76}	72.00 _{4.44}	82.63 _{1.80}	83.07 _{1.56}	83.53 _{1.02}
RTE	50.54 _{3.19}	53.99 _{0.68}	54.13 _{1.19}	49.58 _{4.85}	56.20 _{1.80}	52.37 _{0.34}	52.91 _{0.45}	53.15 _{0.29}	53.30 _{1.28}	56.16 _{1.23}
Avg.	65.60	65.78	67.17	66.61	69.34	51.65	60.69	64.48	65.32	66.32

Table 2: Additional experiments on streamlining grid search by adopting a fixed set of thresholds. The **best** and second-best results are marked by bold and undrline, respectively. Experiments are conducted using Llama2-7b.

	SST2	MMLU	COPA	RTE	MR	Trec	Avg.
ICL	87.22 _{6.03}	41.73 _{2.25}	67.60 _{2.30}	66.21 _{7.30}	89.37 _{1.83}	72.92 _{12.42}	70.84
CC	92.24 _{3.39}	43.72 _{0.97}	67.80 _{2.17}	64.33 _{3.68}	91.77 _{1.42}	76.44 _{3.21}	72.72
DC	94.15 _{1.22}	43.57 _{1.38}	60.40 _{2.79}	65.49 _{2.09}	92.35 _{0.23}	77.16 _{3.94}	72.19
PC	93.90 _{1.54}	34.12 _{3.41}	67.80 _{3.70}	62.59 _{4.71}	91.39 _{1.65}	74.92 _{5.78}	70.79
UniBias	94.54 _{0.62}	44.83 _{0.24}	69.00 _{2.74}	67.65 _{6.44}	92.19 _{0.37}	80.80 _{3.17}	74.84
UniBias with fixed thresholds	<u>94.42</u> _{0.80}	<u>44.47</u> _{0.93}	<u>68.20</u> _{1.79}	<u>67.51</u> _{4.97}	92.35 _{0.17}	<u>80.32</u> _{4.40}	<u>74.54</u>

Table 3: Shared biased attention heads and their frequency of occurrence across 12 datasets in the Llama2-7b model. Each entry represents a specific (layer index, head index) combination.

(19, 10)	(19, 14)	(16, 29)	(19, 21)	(25, 21)	(16, 11)	(18, 31)	(18, 1)
6	5	4	4	3	3	3	3

Table 4: Additional experiments on eliminating common biased components. Attention heads list in Table 3 are removed from the original Llama2-7b model, and the modified model is then evaluated across multiple tasks.

	SST2	MMLU	COPA	RTE	MR	Trec	Avg.
ICL	87.22 _{6.03}	41.73 _{2.25}	67.60 _{2.30}	66.21 _{7.30}	89.37 _{1.83}	72.92 _{12.42}	70.84
Unibias	94.54 _{0.62}	44.83 _{0.24}	69.00 _{2.74}	67.65 _{6.44}	92.19 _{0.37}	80.80 _{3.17}	74.84
Eliminating Common Biased Components	94.32 _{0.60}	44.20 _{1.14}	68.00 _{2.87}	67.37 _{4.60}	92.43 _{0.09}	77.60 _{4.75}	73.98

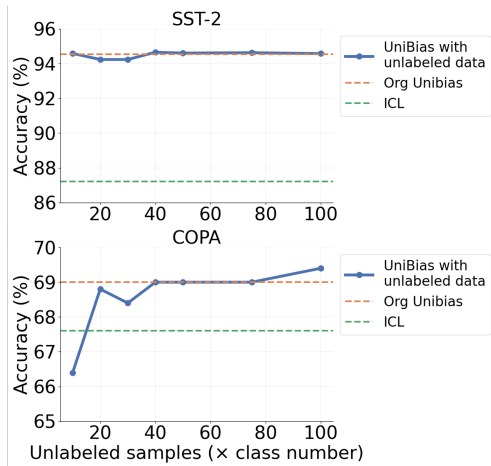


Figure 1: Performance of Unibias using unlabeled samples as support set. It is compared against standard ICL and the original Unibias method.

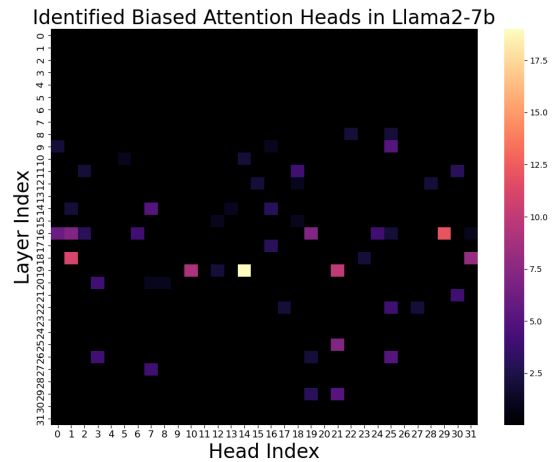


Figure 2: Identified biased attention heads across 12 datasets with 5 repetitions for each dataset.