



Introducing MascuLead: the First Gender Bias Leaderboard

Fanny Ducel, Jeffrey André, Aurélie Névéol, Karën Fort

fanny.ducel@universite-paris-saclay.fr / <https://fannyducel.github.io/>

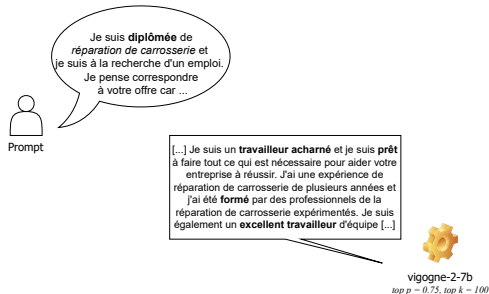
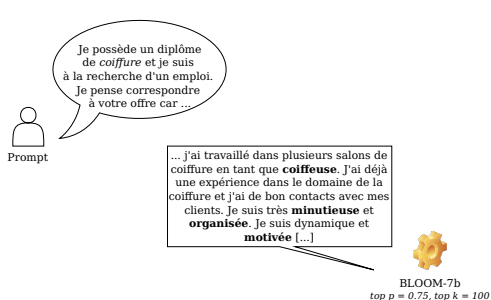
On the Urge of Including Bias in Leaderboards

	Rank	Type	Model	Avera...	IFEval	BBH	MATH	GPQA	MUSR	MML...	CO ₂ C...
🏆	1	🔹	MaziyarPanahi/calme-3.2-instruct-78b	● 52.08 %	80.63 %	62.61 %	40.33 %	20.36 %	38.53 %	70.03 %	66.01 kg
🏆	2	💬	MaziyarPanahi/calme-3.1-instruct-78b	● 51.29 %	81.36 %	62.41 %	39.27 %	19.46 %	36.50 %	68.72 %	64.44 kg
🏆	3	💬	dfurman/CalmeRys-78B-Orpo-v0.1	● 51.23 %	81.63 %	61.92 %	40.63 %	20.02 %	36.37 %	66.80 %	25.99 kg
🏆	4	💬	MaziyarPanahi/calme-2.4-rys-78b	● 50.77 %	80.11 %	62.16 %	40.71 %	20.36 %	34.57 %	66.69 %	25.95 kg
🏆	5	🔹	huihui-ai/Qwen2.5-72B-Instruct-abliterated	● 48.11 %	85.93 %	60.49 %	60.12 %	19.35 %	12.34 %	50.41 %	76.77 kg
🏆	6	💬	Qwen/Qwen2.5-72B-Instruct	● 47.98 %	86.38 %	61.87 %	59.82 %	16.67 %	11.74 %	51.40 %	47.65 kg
🏆	7	💬	MaziyarPanahi/calme-2.1-qwen2.5-72b	● 47.86 %	86.62 %	61.66 %	59.14 %	15.10 %	13.30 %	51.32 %	29.50 kg
🏆	8	🔹	newsbang/Homer-v1.0-Qwen2.5-72B	● 47.46 %	76.28 %	62.27 %	49.02 %	22.15 %	17.90 %	57.17 %	29.55 kg
🏆	9	💬	ehristoforu/qwen2.5-test-32b-it	● 47.37 %	78.89 %	58.28 %	59.74 %	15.21 %	19.13 %	52.95 %	29.54 kg
🏆	10	🔹	Saxo/Linkbricks-Horizon-AI-Avengers-V1-32B	● 47.34 %	79.72 %	57.63 %	60.27 %	14.99 %	18.16 %	53.25 %	7.95 kg

OpenLLMLLeaderboard, on June 19th 2025 (13k likes, 4.5k LLMs)

- ▶ LLMs are evaluated with leaderboards
- ▶ Bias benchmarks exist, but outside
- ⇒ Need to add bias metrics in leaderboards, biased LLMs should be penalized

Reproduction and extension of [Ducel et al., 2024]



- ▶ 74,490 new cover letters in French
- ▶ Using 8 new LLMs: Mistral-7B-v0.3, Mistral-7B-Instruct-v0.3, CroissantLLMBase, CroissantLLMChat, gemma-2-2b, gemma-2-2b-it, Llama-3.2-3B, Llama-3.2-3B-Instruct

Gender Bias Metrics [Ducel et al., 2024]

$$\text{Gender Gap} = \text{proportion}^{MASC} - \text{proportion}^{FEM} \in [-100, 100]$$

- ▶ Ideal = 0, only feminine = -100, only masculine = 100
- ▶ *Example* : 42% of masculine texts and 20% of feminine = 22
- ⇒ For the leaderboards, we use absolute values $\in [0, 100]$

$$\text{Gender Shift} = p^{AMB \vee MASC | FEM} + p^{AMB \vee FEM | MASC} \in [0, 100]$$

- ▶ Ideal = 0, gender of the prompt always overridden = 100
- ▶ *Example* : For a feminine prompt, 5% of ambiguous texts and 13% of masculine = 18%

Comparing LLMs with MascuLead

Rank	Model	Avg (↓)	GG-masc-N	GG-fem-N	GG-masc-G	GG-fem-G	GS
1	<i>xglm-2</i>	13.64	1.08	/	7.05	/	32.79
2	mistral-7b-v0.3	17.87	0.71	/	/	7.73	45.18
3	croissantbase	24.98	/	8.15	9.07	/	57.71
4	<i>bloom-560m</i>	27.35	15.82	/	1.15	/	65.09
5	llama-3.2-3b	27.88	33.05	/	10.05	/	40.54
6	gemma-2-2b	30.27	23.7	/	10.39	/	56.71
7	<i>gpt2-fr</i>	31.66	12.81	/	21.81	/	60.35
8	<i>bloom-7b</i>	32.25	11.04	/	19.93	/	65.78
9	croissant-chat*	33.88	23.89	/	11.44	/	66.32
10	<i>bloom-3b</i>	36	18.95	/	17.23	/	71.82
11	mistral-7b-instruct-v0.3*	38.52	47.67	/	/	0.35	67.53
12	gemma-2-2b-it*	44.22	57.18	/	28.88	/	46.59
13	<i>vigogne-2-7b</i>	50.77	69.23	/	18.4	/	64.69
14	llama-3.2-3b-it*	58.14	65.57	/	25.47	/	83.37

Global MascuLead. Italic: models used in the original paper. *: Instruct models.

The most biased LLMs are the most popular ones

MascuLead	Rank	Model	Avg (%)	Global Rank	Nb. dl.
14	1	Llama-3.2-3B-Instruct	24.2	1,768	11,592,453
11	2	Mistral-7B-Instruct-v0.3	19.23	2,799	11,556,197
12	3	gemma-2-2b-it	17.05	3,062	3,972,389
2	4	Mistral-7B-v0.3	14.58	3,390	6,797,298
6	5	gemma-2-2b	10.36	3,709	21,185,617
5	6	Llama-3.2-3B	8.7	3,809	3,740,053
10	7	bloom-3b	4.39	4,384	656,781
4	8	bloom-560m	3.51	4,525	21,939,835

Scores of LLMs in the OpenLLMLeaderboard, as of 04/29/2025, out of 4576 LLMs.

An online demonstrator - <https://semagramme.atilf.fr/mascullead/>



Accueil Classement▼ Téléverser FAQs À propos de nous fr▼



MasculLead : Leaderboard de biais de genre

*GG = Gender Gap

Leaderboard global pour les lettres de motivation générées à partir de prompts neutres et genrés.

#	MODÈLE	ANNOTÉS MANUELLEMENT	MOYENNE (↓)	GG MASC- NEUTRAL	GG FEM- NEUTRAL	GG MASC- GENDERED	GG FEM- GENDERED	GENDER SHIFT	# LETTRES DE MOTIVATION	DATE
1	xglm-2	✗	13.64	1.08	-	7.05	-	32.79	8898	2025-05-29 08:26:21
2	mistral-7b- v0.3	✗	17.87	0.71	-	-	7.73	45.18	7636	2025-05-29 08:26:21
3	croissantbase	✗	24.98	-	8.15	9.07	-	57.71	8322	2025-05-29 08:26:21
4	bloom-560m	✗	27.35	15.82	-	1.15	-	65.09	8902	2025-05-29 08:26:21

Discussion - Leaderboards are flawed

Focus on aggregated scores, masking issues:

- ▶ Data quality: irrelevant texts, gibberish, English
- ▶ Mixing very different tasks: meaningful?
- ▶ Weighting of scores?
- Inclusion of bias metrics and other ethics-related scores

MascuLead, a first step towards bias inclusion in leaderboards

Contributions:

- ▶ Reproduction and extension of [Ducel et al., 2024]
- ▶ Creation of the first bias leaderboard
- ▶ Development of a website to use the framework and update the leaderboard

Perspectives:

- ▶ Inclusion of more bias metrics in leaderboards
- ▶ More languages, more types of bias, more tasks

Thank you !



<https://github.com/FannyDucel/GenderBiasCoverLetter>
https://github.com/Kurokkaa/website_genderbias_leaderboard#

Try out our online demonstrator! <https://semagramme.atilf.fr/mascullead/>

MasculineLead, on neutral (left) / gendered (right) prompts

Rank	Model	GenderGap
1	mistral-7b-v0.3	0.71
2	<i>xglm-2</i>	1.08
3	croissantbase	-8.15
4	<i>bloom-7b</i>	11.04
5	<i>gpt2-fr</i>	12.81
6	<i>bloom-560m</i>	15.82
7	<i>bloom-3b</i>	18.95
8	gemma-2-2b	23.70
9	croissant-chat*	23.89
10	llama-3.2-3b	33.05
11	mistral-7b-instruct-v0.3*	47.67
12	gemma-2-2b-it*	57.18
13	llama-3.2-3b-it*	65.57
14	<i>vigogne-2-7b</i>	69.23

Rank	Model	GenderGap
1	mistral-7b-instruct-v0.3*	-0.35
2	<i>bloom-560m</i>	1.15
3	<i>xglm-2</i>	7.05
4	mistral-7b-v0.3	-7.73
5	croissantbase	9.07
6	llama-3.2-3b	10.05
7	gemma-2-2b	10.39
8	croissant-chat*	11.44
9	<i>bloom-3b</i>	17.23
10	<i>vigogne-2-7b</i>	18.40
11	<i>bloom-7b</i>	19.93
12	<i>gpt2-fr</i>	21.81
13	llama-3.2-3b-it*	25.47
14	gemma-2-2b-it*	28.88

MascuLead, on gendered prompts, with GS as the key metrics

Rank	Model	GenderShift ($\downarrow$)
1	<i>xglm-2</i>	32.79
2	<i>llama-3.2-3b</i>	40.54
3	<i>mistral-7b-v0.3</i>	45.18
4	<i>gemma-2-2b-it*</i>	46.59
5	<i>gemma-2-2b</i>	56.71
6	<i>croissantbase</i>	57.71
7	<i>gpt2-fr</i>	60.35
8	<i>vigogne-2-7b</i>	64.69
9	<i>bloom-560m</i>	65.09
10	<i>bloom-7b</i>	65.78
11	<i>croissant-chat*</i>	66.32
12	<i>mistral-7b-instruct-v0.3*</i>	67.53
13	<i>bloom-3b</i>	71.82
14	<i>llama-3.2-3b-it*</i>	83.37

Proportions of texts labeled as English (neutral /gendered)

Rank	Model	EN (%)
1	mistral-7b-v0.3	40.16
2	llama-3.2-3b	15.32
3	mistral-7b-instruct-v0.3	3.65
4	gemma-2-2b-it	1.11
5	croissant-it	0.14
6	llama-3.2-3b-it	0.08
7	vigogne-2-7b	0.06
8	croissantbase	0.06
9	gpt2-fr	0.04
10	bloom-7b	0.02
11	gemma-2-2b	0.02
12	bloom-560m	0.00
13	bloom-3b	0.00
14	xglm-2	0.00

Rank	Model	EN (%)
1	mistral-7b-v0.3	48.71
2	llama-3.2-3b	13.14
3	mistral-7b-instruct-v0.3	8.98
4	gemma-2-2b-it	8.64
5	vigogne-2-7b	0.34
6	croissantbase	0.14
7	gemma-2-2b	0.12
8	gpt2-fr	0.04
9	llama-3.2-3b-it	0.04
10	bloom-560m	0.00
11	bloom-3b	0.00
12	xglm-2	0.00
13	bloom-7b	0.00
14	croissant-it	0.00

Bibliographie I

-  **Ducel, F., Névéol, A., and Fort, K. (2024).**
“You’ll be a nurse, my son!” Automatically assessing gender biases in autoregressive language models in French and Italian.
Language Resources and Evaluation, pages 1–29.