

CORRECTION

# Correction: GATmath and GATLc: Comprehensive benchmarks for evaluating Arabic large language models

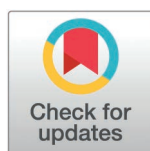
The *PLOS One* Staff

## Notice of Republication

This article was republished on November 14, 2025, to correct errors in the Data Availability Statement. The publisher apologizes for the errors. Please download this article again to view the correct version. The originally published, uncorrected article and the republished, corrected articles are provided here for reference.

## Reference

1. AlBallaa S, AlTwaresh N, AlSalman A, Alfarhood S. GATmath and GATLc: Comprehensive benchmarks for evaluating Arabic large language models. PLoS One. 2025;20(9):e0329129. <https://doi.org/10.1371/journal.pone.0329129> PMID: [40892946](https://pubmed.ncbi.nlm.nih.gov/40892946/)



## OPEN ACCESS

**Citation:** The *PLOS One* Staff (2025) Correction: GATmath and GATLc: Comprehensive benchmarks for evaluating Arabic large language models. PLoS One 20(12): e0338465. <https://doi.org/10.1371/journal.pone.0338465>

**Published:** December 9, 2025

**Copyright:** © 2025 The *PLOS One* Staff. This is an open access article distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.