

Evaluating, Ranking, and Testing Ground Motion Models

Matteo Taroni, Aybige Akinci, Paolo Zimmaro ,
Gianluca Regina, Giovanni Lanzano, Jacopo Selva

Performance of the Models

Classical Approach: Scherbaum et al. 2009 BSSA

$$\langle \log_b(\mathcal{L}(g|\mathbf{x})) \rangle = \frac{1}{N} \sum_{i=1}^N \log_b(g(x_i)).$$



Performance of the Models

Weighted Likelihood Approach: inspired by Zhuang 2015 Spatial Statistics

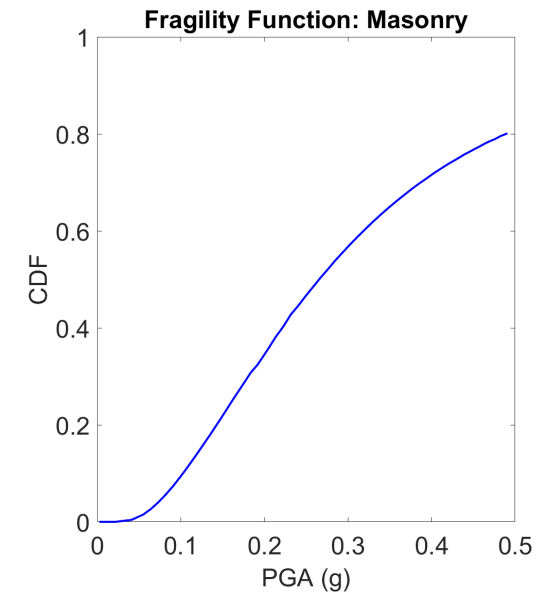
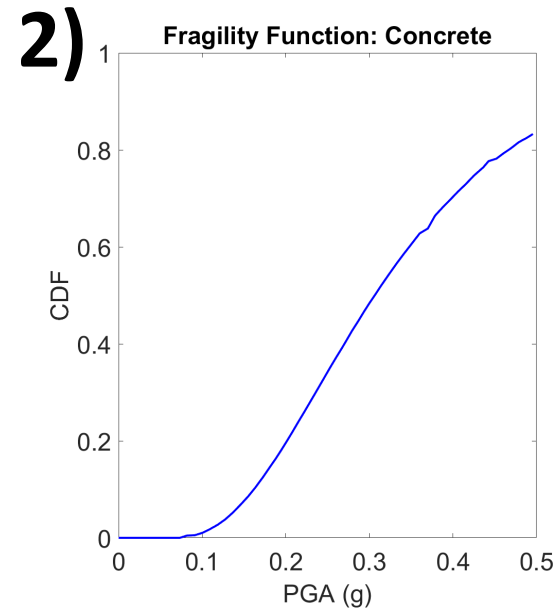
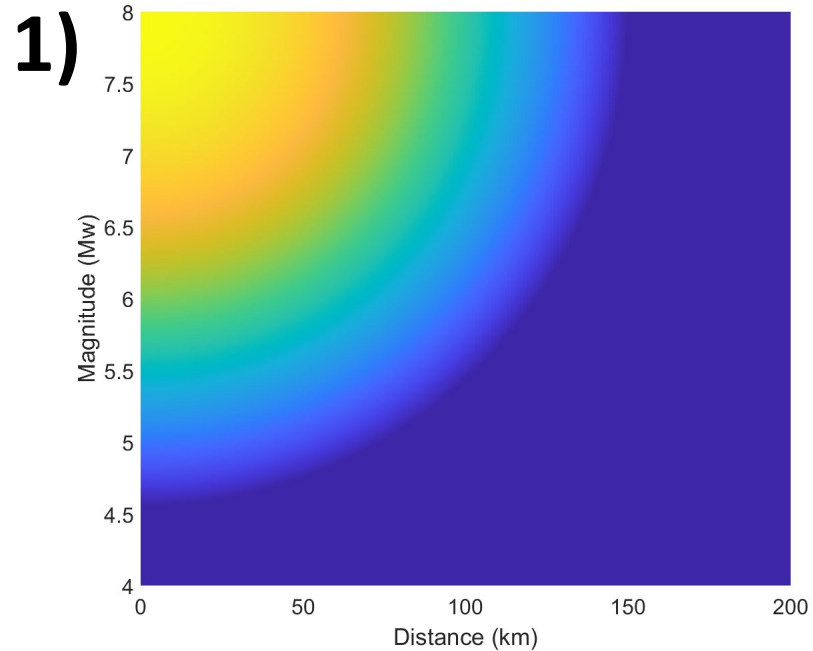
$$\langle \log_b(\mathcal{L}(g|\mathbf{x})) \rangle = \frac{1}{N} \sum_{i=1}^N \log_b(g(x_i)).$$


$$\langle \log_b(\mathcal{L}(g|\mathbf{x})) \rangle = \sum_{i=1}^N \mathbf{w}_i \log_b(g(x_i))$$

$w_i \rightarrow$ *this weight can be used to adapt the score*

Performance of the Models

w_i proportional to ($\approx$) ...



3) $w_i \approx 1/N_k$

Performance of the Models

Table 1: LLH and WLH computed for the different GMMs using the Central Italy observations; WLH (Con) is the WLH computed using the weights proportional to the concrete buildings fragility function, while WLH (Mas) is the WLH for masonry buildings.

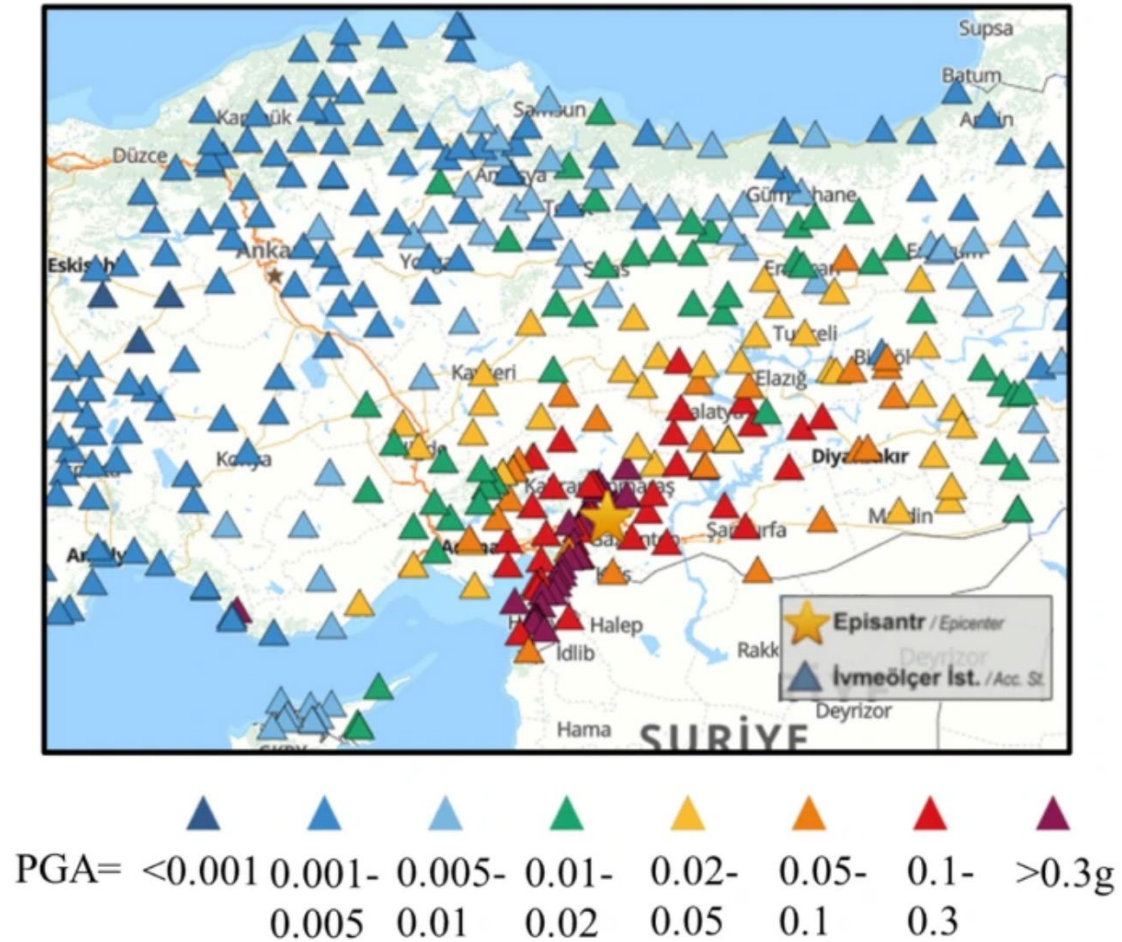
Model	LLH	WLH (Con)	WLH (Mas)
ITA10	1.86	2.28	2.38
BIN14	1.98	2.02	2.13
BSSA14	1.94	2.47	2.66
ITA18	1.79	2.29	2.35
KOT22	2.04	1.84	1.94

Performance of the Models

Table 2: rank of the GMMs using the previously computed LLH, WLH (Con), and WLH (Mas).

Ranking LLH	Ranking WLH (Con)	Ranking WLH (Mas)
1) ITA18	1) KOT22	1) KOT22
2) ITA10	2) BIN14	2) BIN14
3) BSSA14	3) ITA10	3) ITA18
4) BIN14	4) ITA18	4) ITA10
5) KOT22	5) BSSA14	5) BSSA14

Testing Models vs Observations



Map of the recording stations of the Kahramanmaraş earthquake during first event (Adapted from AFAD)

Lashgari et al. 2023

Testing Models vs Observations

Using the log-likelihood approach as a distance between model and observations



Our goal is not to compute the real probability of the observations given the model; we rather want to assign a “score” to each output of the model

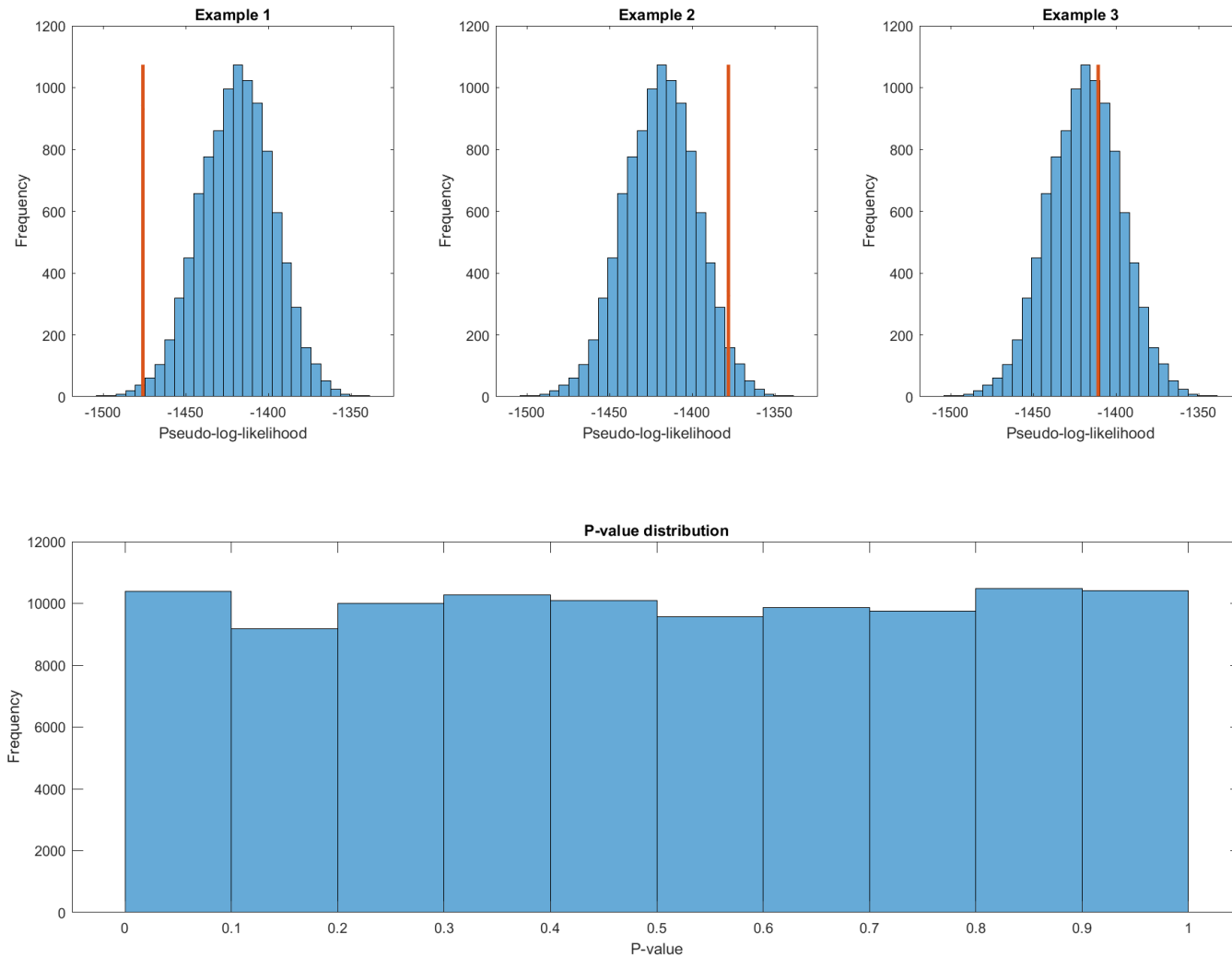
Testing Models vs Observations

Monte Carlo approach: from the model we can simulate thousands of scenarios, for each scenario we can compute the score of the model vs the simulated observations.

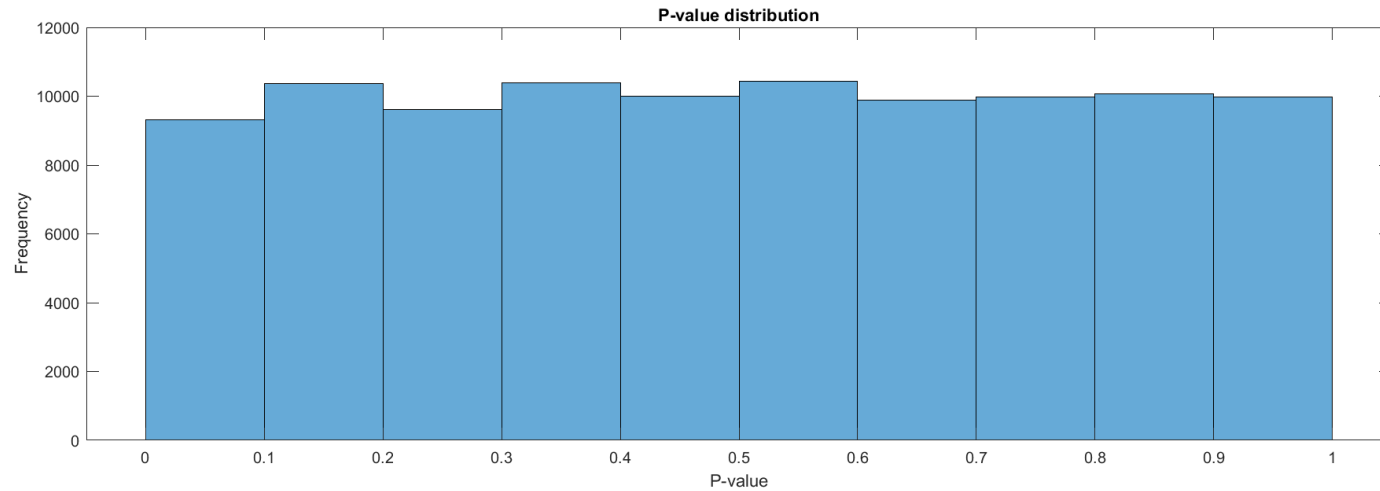
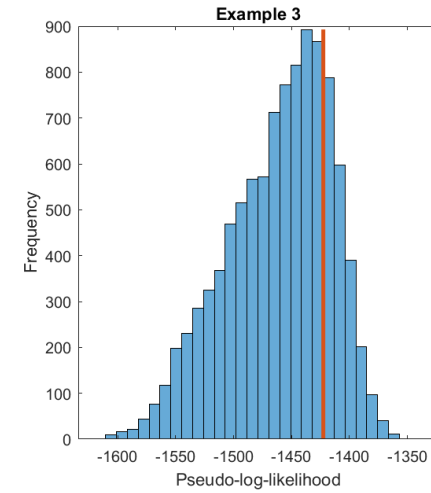
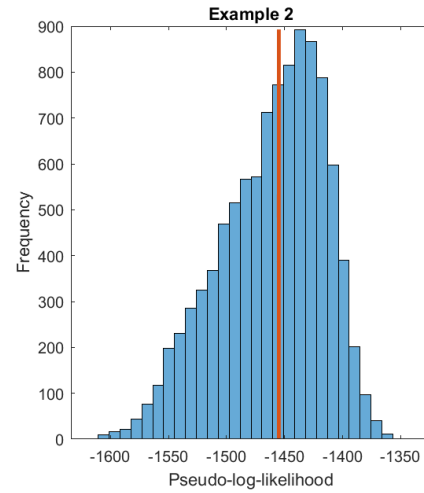
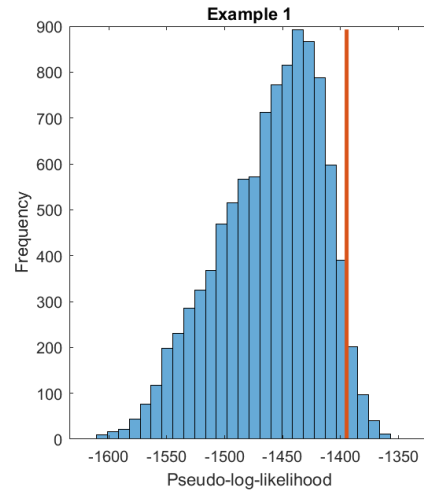


In this manner, we can compute our test statistic distribution; then we have just to compute the score of the model vs the real observations

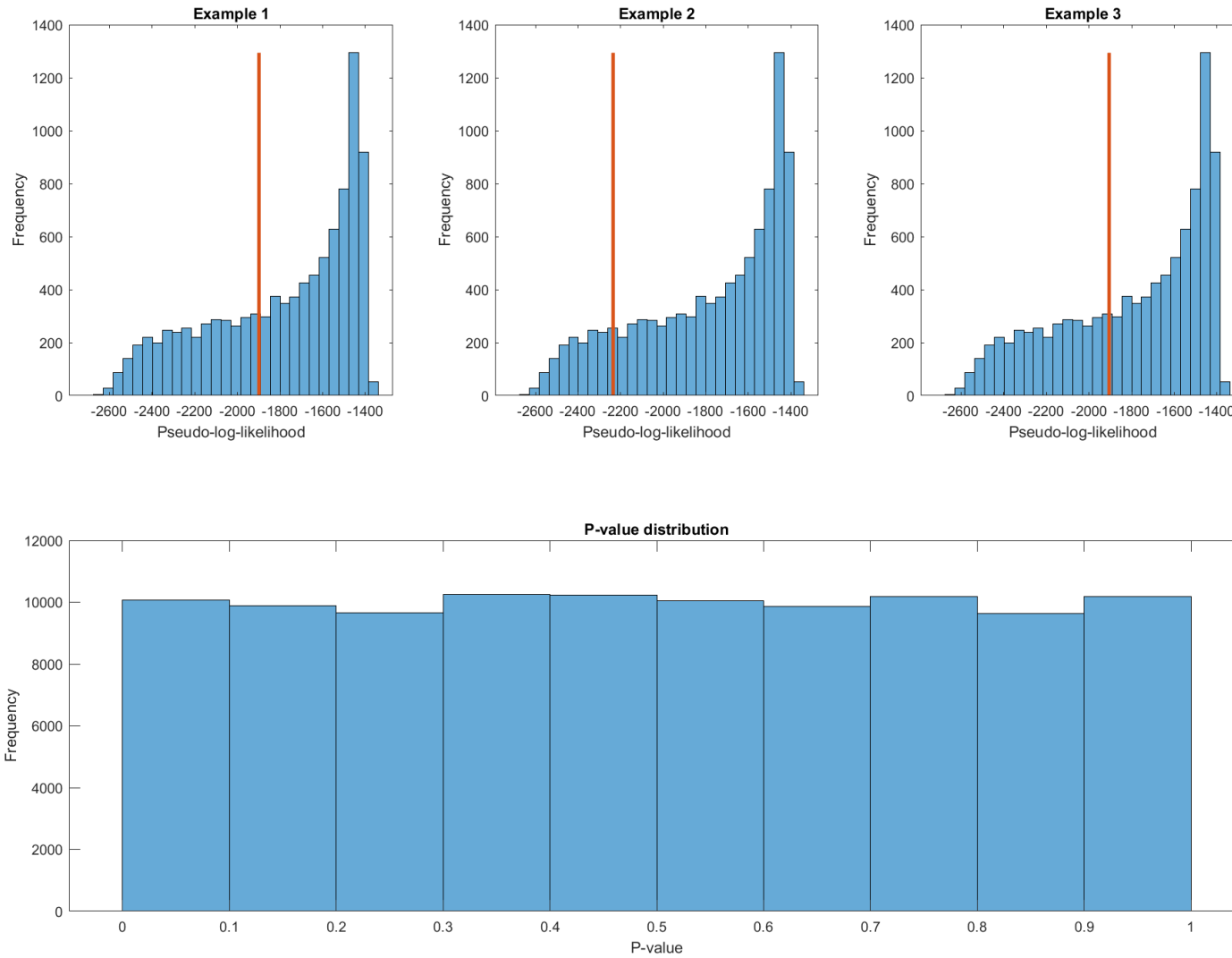
Testing Models vs Observations



Testing Models vs Observations



Testing Models vs Observations





The End