

1. Conveying accurate confidence is critical



Chest X-ray Report

Impression:

- Possible pneumonia. **Order Followup CT**
- Mild bronchial wall thickening, consistent with bronchitis.
- No acute abnormalities in the left lung or mediastinum.

Chest X-ray Report

Impression:

- Likely pneumonia. **Prescribing Antibiotics & Order CT**
- Mild bronchial wall thickening, consistent with bronchitis.
- No acute abnormalities in the left lung or mediastinum.

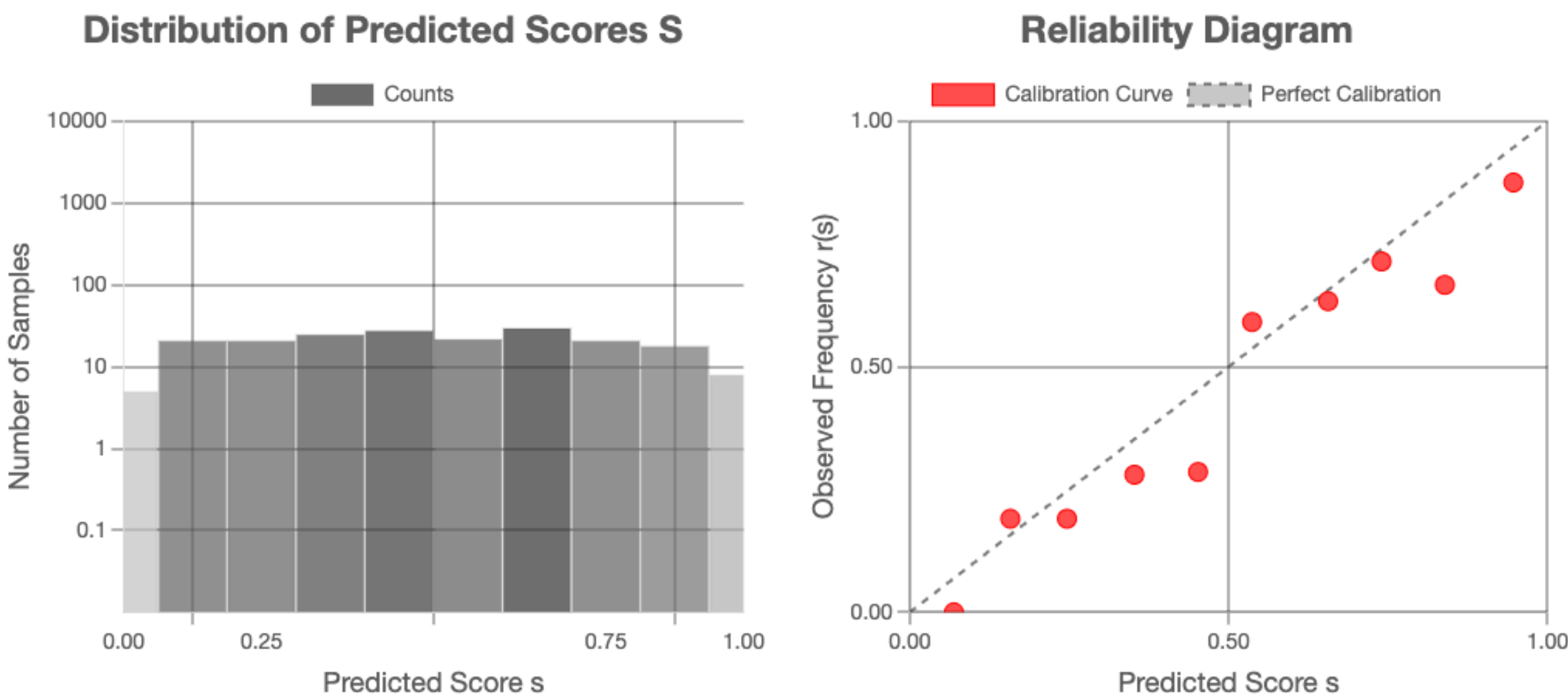
2. Background on Calibration and Miscalibration

- Classifier $g : \mathcal{X} \rightarrow [0, 1]$ is **calibrated** if

“Observed Frequency” $\mathbb{E}[Y | S] = S$
 “Predicted Confidence Score”,

 where $\mathcal{X}$ is input space, $\mathcal{Y} = \{0, 1\}$ is label space, and $S = g(X)$ represents the confidence score.
- Expected Calibration Error** measures the degree to which the calibration definition is violated

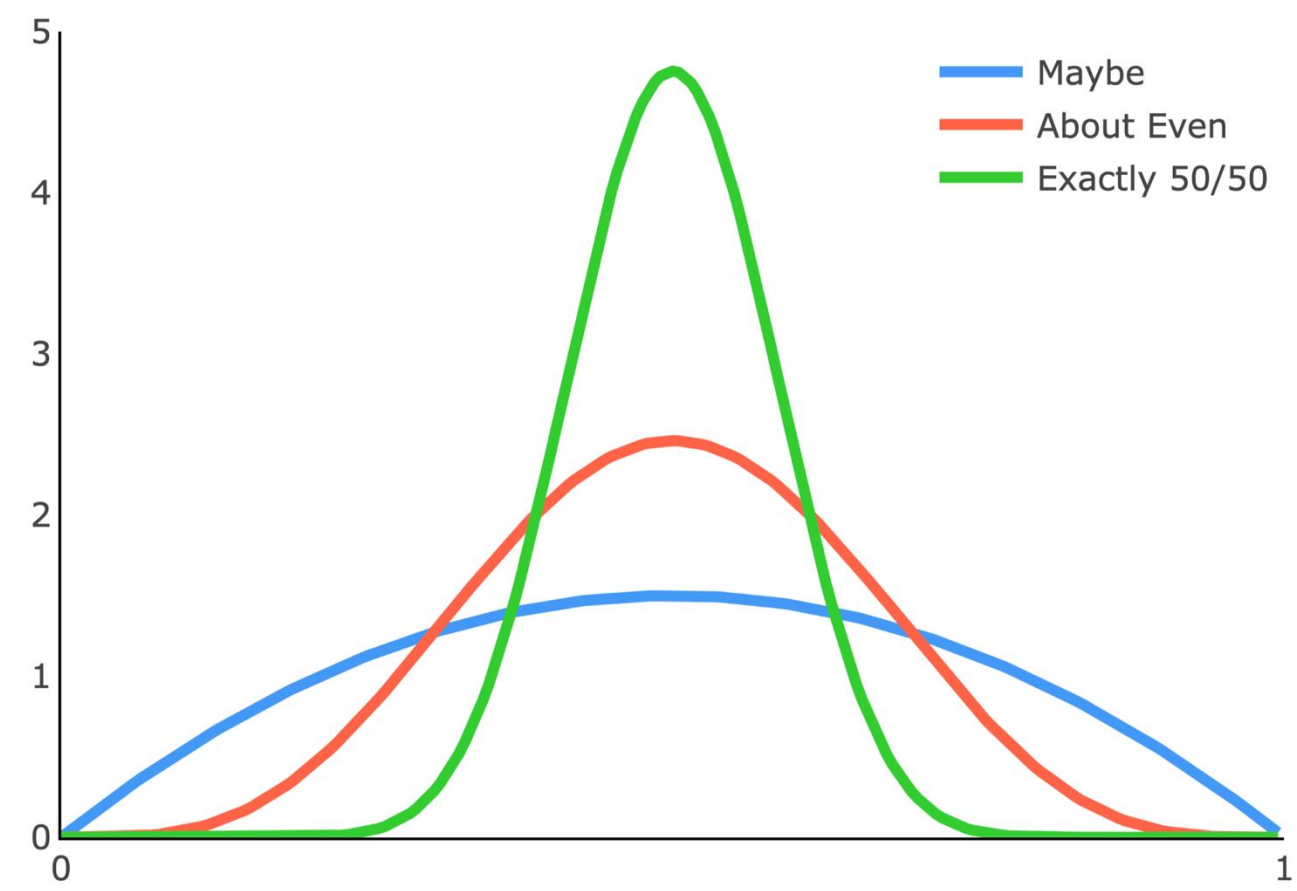
$$ECE = \mathbb{E}[|\mathbb{E}[Y | S] - S|]$$



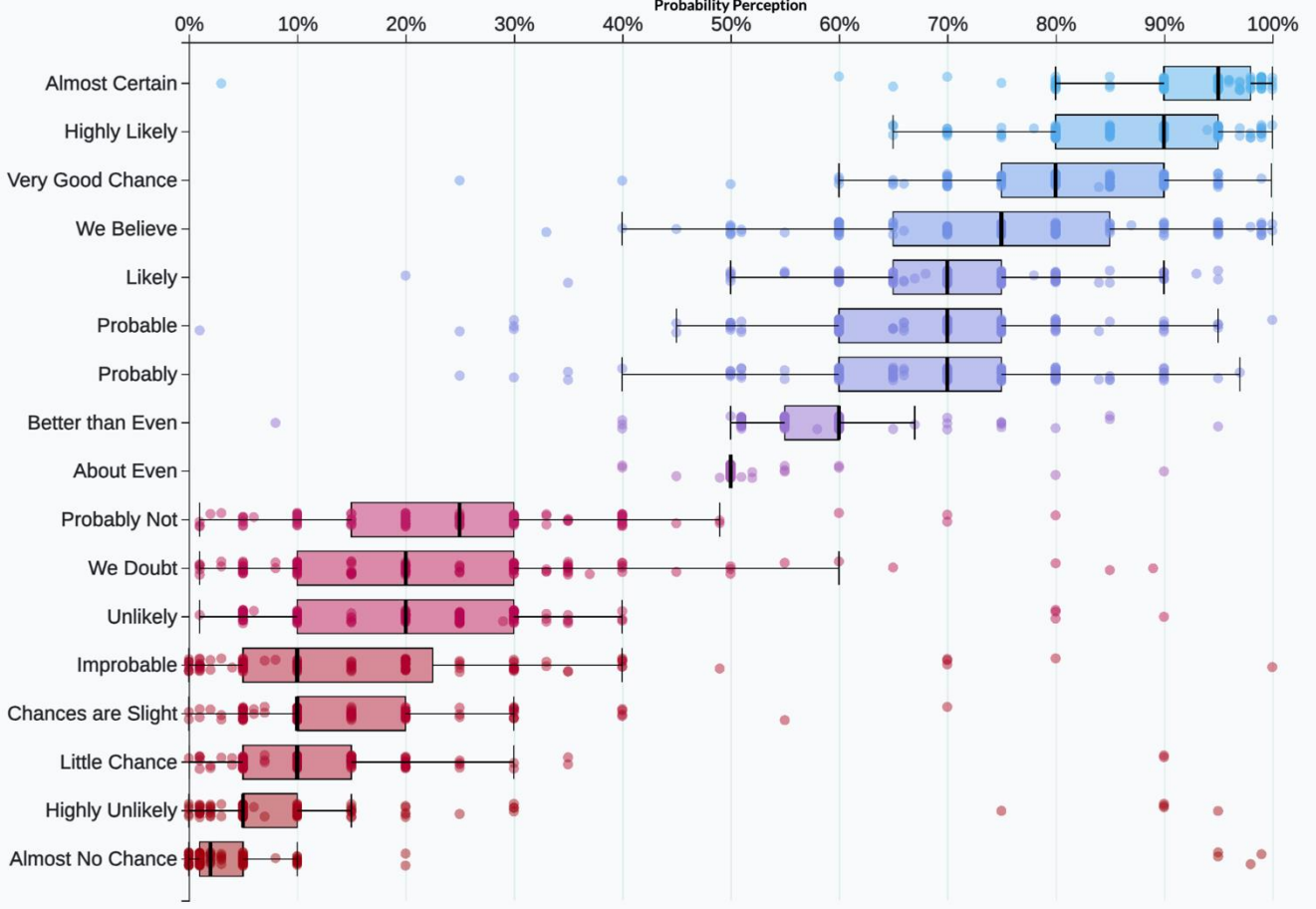
3. Linguistic Expressions of Certainty, not Scores!

assuming certainty phrase $\leftrightarrow$ probability (e.g., “Maybe” $\leftrightarrow$ 0.5, “Likely” $\leftrightarrow$ 0.8, ...) is insufficient

a) Nuanced semantics

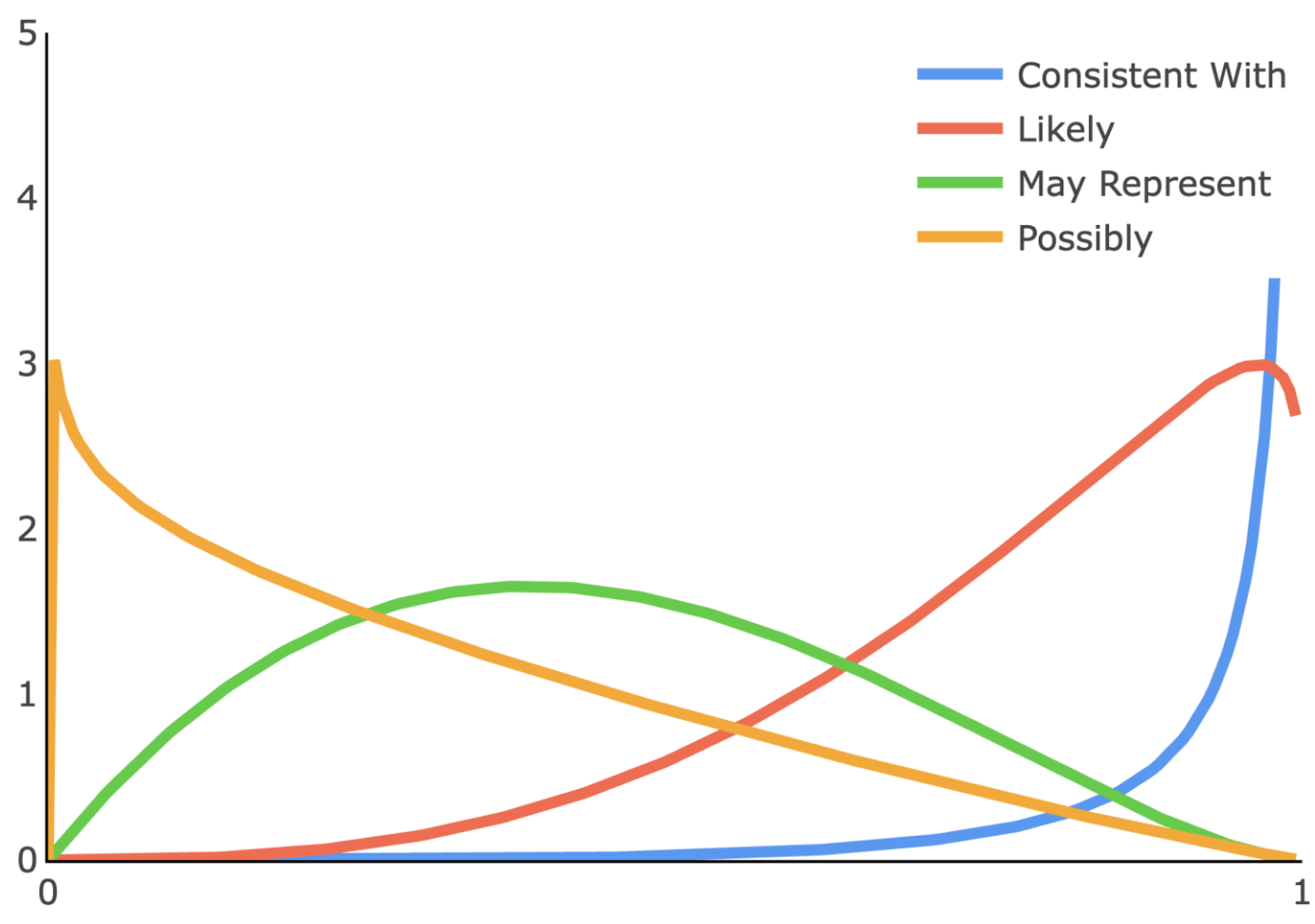


b) Population-level variation



4. Confidence as a Distribution

- assume certainty phrase “Maybe” $\leftrightarrow$ Beta(2,2), “Likely” $\leftrightarrow$ Beta(4,1), ...
- distributions can be derived from (1) survey of human perception of these phrases, (2) prompting LMs (e.g., gpt-4)



6. Measuring Radiologists Calibration

2,662 paired (X-ray, CT) report where X-ray are human predictions and CT are ground-truth

Chest X-ray Report

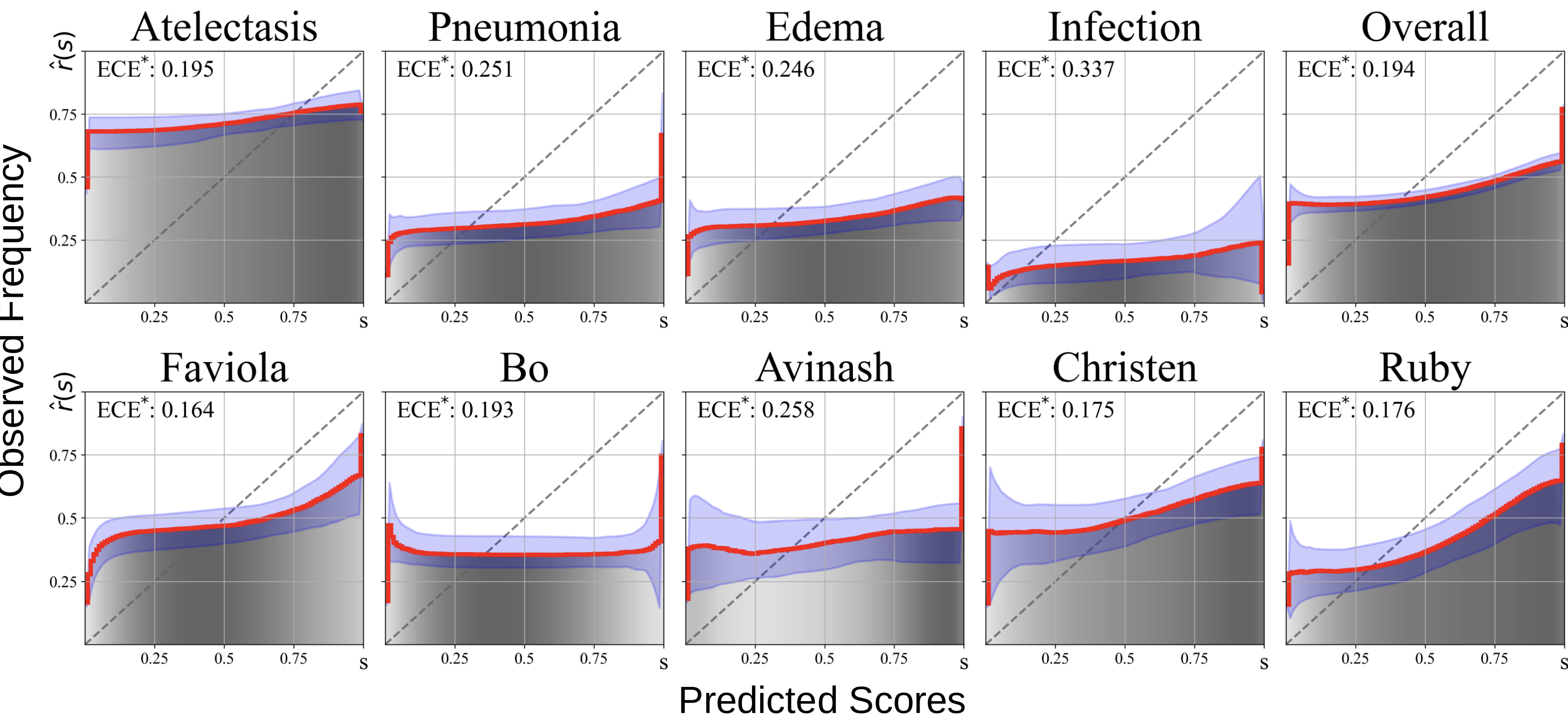
Impression:

- Right lower lobe opacification, which **may represent** pneumonia.
- No acute abnormalities in the left lung or mediastinum.
- Possible pulmonary embolism.

CT Report

Impression:

- Right lower lobe opacification **consistent with** pneumonia.
- No acute abnormalities in the left lung or mediastinum.
- pulmonary embolism is **absent**.



The calibration curve (red), with its 95% confidence interval (blue), and score density (gray) are shown.

7. Post-hoc Calibration

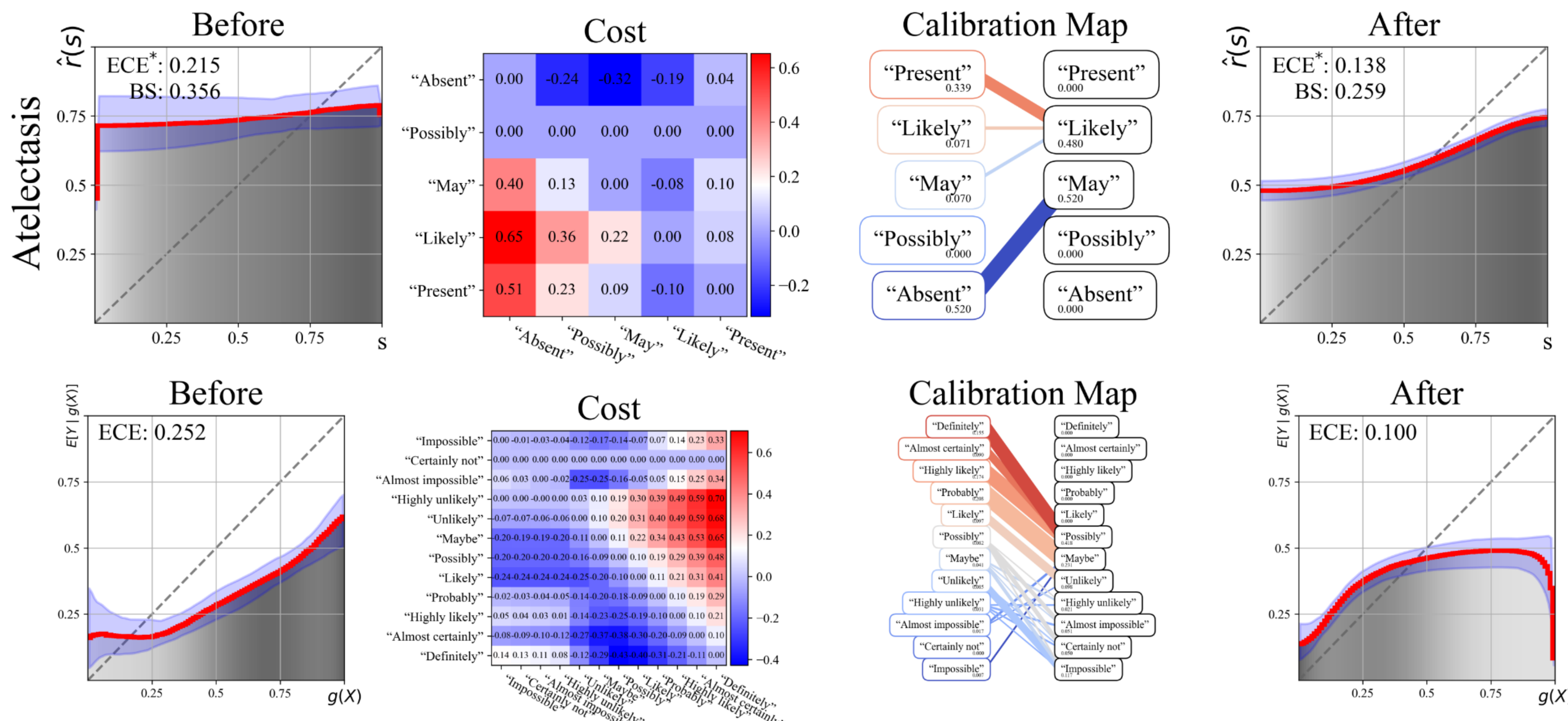
by adjusting the usage of certainty phrases

- Traditional Calibration Map: $t : [0, 1] \rightarrow [0, 1]$ (e.g., platt scaling, histogram binning).
- Our work: $t : [K] \rightarrow [K]$ mapping K source distributions to K target distributions.



Paper Link

8. Calibrating Radiologists (top) and Language Models (bottom)



- Source distribution: $a = (a_1, \dots, a_K) \in \Delta^{K-1}$, where a_k is the proportion of times the k -th certainty phrase is used.
- Cost matrix $C \in \mathbb{R}^{K \times K}$: $C_{kl} = \frac{1}{a_k} (\widehat{ECE}(u_k \rightarrow v_l) - \widehat{ECE})$.
- Calibration as discrete OT: $\min_T \langle C, T \rangle$ s.t. $T\mathbf{1} = a$ (preserve source distribution), $T^\top \mathbf{1} \approx a$ (target close to source).

This work was supported by the Takeda Fellowship, the MIT–IBM Watson AI Lab, the MIT CSAIL–Wistron Program, and the MIT Jameel Clinic.