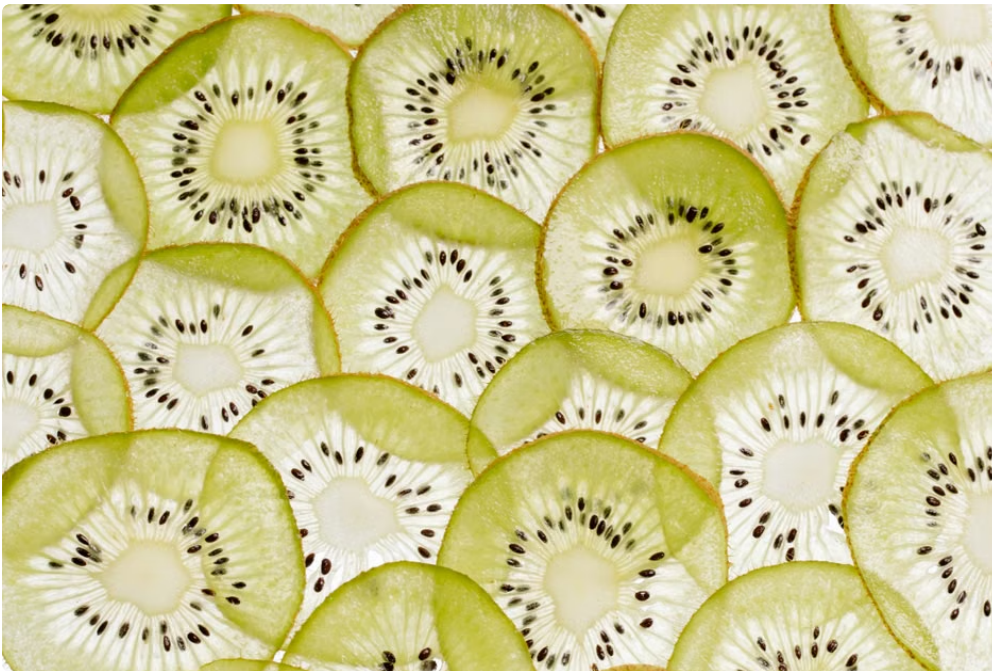


When teams tackle the perennial challenge of labeling disputed instances in machine learning, a common heuristic surfaces: “Label at least 500 disputed instances.” But as a 12-year practitioner who has shipped risk-scored decision systems in lending and healthcare — areas where mistakes cost real lives and money — I find that simply hitting a round number without context can be dangerously misleading. Let’s deeply unpack what underlies this rule <https://reportz.io/ai/when-models-disagree-what-contradictions-reveal-that-a-single-ai-would-miss/> of thumb through the lenses of **disagreement rate**, **predictive entropy**, and other considerations.

Why Focus on Disputed Instances?

Disputed instances, or data points the model struggles to confidently categorize, represent the “edge cases” that define the limits of your model’s understanding. They live in what I call *disagreement regions*: areas where the model’s predicted probabilities are conflicted or unstable, or where multiple models in an ensemble don’t agree.

Consider these disputed instances as the highest-signal risks in your dataset. If you only label the “easy” cases, you won’t learn how your model behaves under real-world stress or distribution shifts. Yet, these are exactly the examples where your model’s performance directly impacts fairness, accuracy, and trust.



Disagreement Rate and Its Role

The **disagreement rate** captures how often two or more predictive models (or versions of the same model via bootstrapping or ensembling) disagree on the label of a given instance. This is an important metric because it tells you where your model’s decisions have the highest uncertainty stemming from competing hypotheses in the learned function.

- **High disagreement rate** zones often align with complex subgroups or rare cases in the data.
- Labeling in these regions is invaluable since it helps reduce uncertainty near decision boundaries.
- Disagreement rate can be used to target sampling for human annotation, making labeling budgets more efficient.

Predictive Entropy: Quantifying Uncertainty from the Model’s POV

Predictive entropy measures uncertainty directly from the model's output probability distribution for a class. High entropy implies the model is "on the fence" — say for binary classification, prediction probabilities near 0.5 generate maximum entropy.

Combined with disagreement rate, predictive entropy helps you identify disputed instances not only from the disagreements among models but also from the model's own uncertainty, providing a multi-dimensional view of uncertainty that is more informative than raw accuracy metrics.

Is Labeling 500 Disputed Instances Enough? The Short Answer: It Depends

We need to move beyond "magic numbers" toward *sample size planning* grounded on actual data distribution, task complexity, and labeling budget constraints. Labeling disputed instances is a critical ingredient, but 500 is arbitrary without considering these factors:

1. **Disagreement region size and coverage:** How many disputed instances exist in your dataset? Are these 500 samples representative of all disagreement regions?
2. **Subgroup and edge case diversity:** Does the sample cover multiple subgroups or rare but high-impact edge cases?
3. **Labeling budget and cost:** Is 500 the maximum you can afford? Do you have prioritized labeling to maximize impact?
4. **Objective mismatch and loss tradeoffs:** Are you optimizing metrics aligned with your use case, or just relying on default losses such as cross-entropy without considering cost-sensitive risks?

Distribution Shift and Edge Cases

Many performance drops happen on rare or shifting subpopulations—the *long tail*—that standard training sets or test splits may not capture. Labeling disputed instances helps uncover these blind spots, but 500 labels may still be insufficient if:

- The underlying distribution shift introduces many new edge cases.
- Your disagreeing instances cluster heavily around only a few subgroups.
- Many disputed samples are correlated or redundant rather than diverse.

When operating under distribution shift, anticipate that you likely need a dynamic labeling approach with ongoing data collection beyond fixed-size batches.

Data Gaps and Subgroup Coverage

One of the hidden traps that accuracy conceals — part of my "things accuracy hides" list — is poor subgroup representation. Disputed instances, especially from disagreement regions, are prime opportunities to discover these gaps.

But if your 500 instances disproportionately represent majority subgroups or similar contexts, your labels will not improve fairness or robustness broadly. Effective sample size planning should include:

- Stratification to ensure some minimal coverage for protected or at-risk groups.
- Augmented strategies such as oversampling or targeted labeling for minority categories.
- Explicit bias or fairness audits informed by disagreement metrics.

Objective Mismatch and Loss Function Tradeoffs

In many projects, teams optimize standard loss functions (e.g., cross-entropy) which may not correspond well with real-world costs or risks. This *objective mismatch* can sow confusion in interpreting disagreement and predictive entropy signals.

- Disagreement regions might include instances where the model's loss gradient is low but cost of misclassification is high.
- By aligning your loss and uncertainty measures with domain-specific costs, you can prioritize labeling disputed samples that reduce your most impactful errors.
- Thresholds for labeling should be tied explicitly to expected utility gains, not just model uncertainty or “vibes.”

In fields like lending or healthcare, misclassifications differ drastically in cost and risk; your labeling efforts must reflect those cost asymmetries.

Takeaways: Designing Your Labeling Budget and Sample Size

Consideration Key Question Impact on Label Sample Size Disagreement Region Size How large and diverse are regions of model disagreement? Large, complex regions require larger labeled sets; targeting 500 may be insufficient. Subgroup and Edge Case Coverage Do disputed instances represent all important subgroups? Sampling must be stratified, possibly increasing labeled samples to cover diversity. Labeling Budget What is your cost and time constraint? Define priorities to maximize ROI on limited sample sizes. Cost-sensitive Objective Are you optimizing for real-world cost, or just accuracy? May require labeling more specific disputed instances aligned with domain risks. Distribution Shift Prediction Are you expecting data drift or new use cases? Prepare an ongoing labeling plan with incremental sample sizes beyond initial 500.

Final Thoughts: Always Ask “What Happens on the Worst Day in Prod?”

Labeling disputed instances is not about ticking a box or hitting “at least 500” labels. It’s about organizing your human-in-the-loop efforts where they are most impactful: reducing your model’s blind spots and risks under real-world use.



Calibration of uncertainty and disagreement metrics, thoughtful sample size planning, and alignment with cost-sensitive objectives must guide your labeling budget. Only then can those 500, or 1000, or 10,000 labels transform your system's trustworthiness—especially when the stakes are high.

And remember, never settle for uncalibrated probabilities or hand-wavy assurances that “AI will handle it.” Ground your labeling strategy in rigorous metrics like *disagreement rate* and *predictive entropy*, tie thresholds explicitly to cost-benefit tradeoffs, and continuously monitor performance—even when the noise in production peaks. That's the only way to ensure your “worst day in prod” isn't your model's last.