

In high-stakes machine learning systems — such as lending decisions, healthcare operations, or fraud detection — automated models often operate alongside human experts to help catch edge cases, mitigate risks from distribution shifts, and address subtle data gaps. One of the most effective ways to formalize this collaboration is by defining a **human escalation policy** based on statistically grounded *disagreement thresholds*. By selecting when a case should be flagged for **human review**, organizations can balance automation scale with risk mitigation, while optimizing for the **cost of errors**.

This post explores the key concepts and tools you should know to set these triage thresholds confidently. We'll focus on two primary signals of uncertainty and disagreement: the **disagreement rate** between models or annotators, and the **predictive entropy** of probabilistic outputs. We'll also discuss how to interpret these signals as high-signal risk indicators that relate to *objective mismatches*, *loss function tradeoffs*, and *subgroup coverage gaps*.

Why Disagreement Is a High-Signal Risk Indicator

At the heart of a **human escalation policy** is the assumption that when automated systems express uncertainty, or when multiple models or annotation sources disagree, there's an elevated chance that an error could be costly if left unreviewed.

- **Disagreement suggests risk:** When two or more models trained on slightly different data, architectures, or objectives disagree, it signals that the example sits near a decision boundary or may be an outlier under the training distribution.
- **Calibration helps, but isn't enough:** Predictive probabilities alone can sometimes be overconfident or poorly calibrated. Disagreement can reveal uncertainty missed by simple confidence scores.
- **Triggers human attention efficiently:** Rather than reviewing all cases or those low-confidence by a single model, triaging based on disagreement maximizes human reviewer impact by catching subtler errors.

But how do you measure and quantify this disagreement? Two common metrics are the *disagreement rate* and *predictive entropy*.

Measuring Disagreement: Disagreement Rate and Predictive Entropy

Disagreement Rate

The disagreement rate measures the proportion of cases where two or more decision-makers—such as models, annotators, or different versions of a model—disagree on the predicted label.

Set A Set B Disagreement? Approve Approve No Deny Approve Yes Deny Deny No

In practice, when multiple models produce probability scores rather than hard classifications, disagreement can be operationalized as thresholded label disagreements—i.e., the models assign different labels after applying a classification threshold—or via more continuous metrics like Kullback–Leibler divergence among predictions.

Predictive Entropy

Predictive entropy is a well-known uncertainty measure derived from the probability distribution over class labels output by a model. It quantifies the *amount of uncertainty* in the prediction itself:

$$\text{Entropy}(p) = -\sum_c p(c) \log p(c)$$

For example, consider [radiology ai human in loop](#) a binary setting where the model outputs probabilities p of the positive class. The entropy is highest when $p = 0.5$ (maximum uncertainty), and lowest when p is near 0 or 1 (confident prediction).

Predictive entropy captures uncertainty arising from overlapping classes or ambiguous data points, which often coincide with cases where the model may make mistakes or where human judgment is valuable.

Setting a Disagreement Threshold for Human Review

Once you have operationalized disagreement or predictive entropy, the key question is how to select a threshold that decides when to escalate to humans. This threshold is critical for balancing the benefits of automation against the risks and costs of errors.

1. Define Your Objectives and Constraints

Before selecting a threshold, clarify the business costs of different error types. For example:

- **False negatives:** Cases where the model incorrectly approves a bad loan or misses a healthcare risk.
- **False positives:** Cases where the model unnecessarily escalates an easy decision, wasting human review capacity.

Understanding these costs helps you frame the threshold choice as an explicit tradeoff between automation efficiency and risk mitigation, avoiding vague “confidence” cutoffs.

2. Analyze Disagreement Distribution Over Historical Data

Collect historical cases with labels verified by experts and compare model predictions. Plot the distribution of disagreement rates or predictive entropy values for:

- Correct predictions
- Incorrect predictions
- Edge-case or flagged examples

This lets you observe whether higher disagreement scores correspond to higher misclassification risk, which is the hallmark of a reliable escalation signal.

3. Quantify Tradeoffs Through Cost-Weighted Metrics

Use metrics that incorporate the **cost of errors**, such as cost-weighted precision, recall, or expected risk. You can simulate or backtest varying disagreement thresholds, then plot:



- Human review volume (how many cases get escalated)
- Expected cost reduction (fewer costly mistakes after review)

This yields a practical curve where you can pick a threshold achieving acceptable risk reduction while controlling review load.

4. Consider Distribution Shift and Data Gaps

Disagreement thresholds set on historical data may perform differently in production where:

- **Distribution shifts:** The characteristics of incoming examples change over time (e.g., new types of loan applications or patient conditions).
- **Subgroup coverage gaps:** Certain protected or minority groups may have unseen patterns causing higher model uncertainty or bias.

Regularly monitoring disagreement rates and human review outcomes across these subgroups can inform dynamic threshold updates or targeted retraining.

5. Integrate Objective Mismatch and Loss Function Tradeoffs

ML models often optimize surrogate losses not perfectly aligned with final business objectives. Disagreement signals often reflect such *objective mismatches*. For example, a model trained to minimize overall error might fail on rare but critical cases.

By systematically escalating high-disagreement cases, you safeguard against such failures, effectively combining automated recall with human judgment to cover the blind spots of your loss function.



Case Study: Implementing a Human Escalation Policy in Lending

In lending, approving a high-risk application incorrectly can cost thousands, whereas over-reviewing low-risk cases wastes human analyst time and slows down customer experience.

Step-by-Step Approach

1. Train two different models (e.g., gradient boosted trees and a neural network) on the same training data.
2. For each application, compute disagreement: do the models disagree on approve/deny?
3. Calculate predictive entropy from each model's softmax probabilities.
4. Analyze historical cases to confirm: disagreement and entropy strongly correlate with loans that were defaulted but initially approved.
5. Quantify cost impact: false approvals cost \$X, human review costs \$Y per case.
6. Backtest various thresholds to find a cutoff maximizing cost savings while limiting reviews to 5% of applications.
7. Deploy threshold in production; monitor disagreement rates, human review outcomes, and error rates.
8. Regularly recalibrate and refine thresholds in response to emerging data shifts.

Things Accuracy Hides — Why Relying on Accuracy Alone Is Dangerous

It's tempting to set human review thresholds based purely on accuracy or confidence scores, but here's why that's a trap I always watch for:

- **Accuracy ignores high-cost errors:** A classifier with 98% accuracy might still miss the 2% of cases that cause the biggest losses.
- **Calibrated probabilities can still be overconfident:** Overconfident probabilities mislead thresholding decisions and result in missing risky cases.
- **Disagreement exposes uncertainty hidden by accuracy:** Two models both getting 98% accuracy may disagree on critical cases—those deserve human eyes.

To avoid this trap, always combine your accuracy metrics with disagreement and entropy analyses and anchor thresholds to explicit cost tradeoffs.

Summary and Best Practices

- **Triage thresholds should be explicitly tied to business costs and risk tolerance.** Don't pick thresholds based on arbitrary confidence cutoffs or "gut feels."
- **Leverage disagreement rate and predictive entropy as complementary signals.** Disagreement points to boundary or edge cases; predictive entropy quantifies per-example uncertainty.
- **Validate thresholds using cost-weighted metrics and backtesting.** Model historic impacts before production rollout.
- **Continuously monitor for distribution shift and subgroup coverage issues.** Recalibrate thresholds when you detect concept drift or emerging data gaps.
- **Don't let accuracy be the only storytelling metric.** Incorporate error costs and uncertainty signals to craft truly effective human escalation policies.

Setting your human review triage threshold well is a core enabler for robust, safe, and scalable ML-powered decision automation. Disagreement and predictive entropy are powerful tools to guide you—what remains is aligning them with your unique costs, risks, and operational constraints. And always keep asking: *what happens on the worst day in prod if we miss this edge case?*