

Machine Intelligence & Responsible Systems
REVIEW | COMMISSIONED SYNTHESIS

Calibrated Abstention and Escalation in Foundation-Model Systems

Anthony Young - Corresponding author

Abstract

This review examines calibrated abstention and escalation in foundation-model systems. The organizing question is when a model should answer, defer, request evidence, or transfer control to a person. Ten related scholarly sources are synthesized through a decision-centered framework spanning problem definition, mechanism, measurement, evaluation, implementation, and governance. The review does not invent experiments, pooled estimates, or unreported quantitative results. It instead evaluates the strength and transferability of the available evidence, with particular attention to treating a fluent refusal style as evidence of calibrated uncertainty. The resulting framework links technical or empirical performance to explicit use conditions and identifies tests that should precede wider adoption in high-consequence decision support with bounded model authority.

Keywords: large language models, abstention, calibration, uncertainty, human escalation

Synthesis lens	Principal question
Mechanism	the coupling among uncertainty estimation, selective prediction, retrieval, and escalation policy
Evaluation	coverage-risk curves, calibration under shift, and error severity at the decision boundary
Primary risk	treating a fluent refusal style as evidence of calibrated uncertainty

1. Introduction

Research on calibrated abstention and escalation in foundation-model systems sits at the boundary between a promising component and a consequential decision. The core question is when a model should answer, defer, request evidence, or transfer control to a person. Answering it requires more than assembling favorable results. It requires a chain of evidence that makes the problem boundary, mechanism, measurement, comparator, uncertainty, and accountable decision visible to the reader.

This article presents a structured narrative review of ten related publications. Sources were selected for topical relevance, traceable publication metadata, and complementary coverage of technical, empirical, and governance questions. No meta-analytic pooling is attempted because the included studies differ in designs, populations, system boundaries, and outcomes. The purpose is decision-oriented synthesis: to identify what each source can support, where transfer remains uncertain, and which evidence would justify the next stage of use.

The central risk is treating a fluent refusal style as evidence of calibrated uncertainty. A top-line metric or broad policy label can appear decisive while concealing the conditions that produced it. Accordingly, the synthesis separates observation from interpretation and implementation from validation. This separation is especially important when the same system creates benefits for one user or institution while transferring cost, risk, or administrative burden to another.

2. Evidence Base and Problem Boundary

A defensible review begins by defining the unit of analysis. For the present topic, the relevant boundary includes inputs, institutional or technical context, the mechanism producing an output, the person or system acting on that output, and the downstream outcome. Evidence falls outside the boundary when it measures only an intermediate proxy yet is used to claim an end-to-end benefit.

Wen et al. (2025) addresses "Know Your Limits: A Survey of Abstention in Large Language Models." In this synthesis, that work is used as a source on problem definition within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Lin et al. (2023) addresses "Generating with Confidence: Uncertainty Quantification for Black-box Large Language Models." In this synthesis, that work is used as a source on conceptual boundary within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Shorinwa et al. (2025) addresses "A Survey on Uncertainty Quantification of Large Language Models: Taxonomy, Open Research Challenges, and Future Directions." In this synthesis, that work is used as a source on measurement within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the

review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Together, these works provide complementary entry points, but their conclusions should not be flattened into a single ranking. Differences in data generation, setting, temporal horizon, and outcome definition may be substantively meaningful. A review should therefore preserve those differences and state where analogy, rather than direct replication, connects one domain to another.

3. Mechanism and System Design

The operative mechanism is the coupling among uncertainty estimation, selective prediction, retrieval, and escalation policy. This formulation discourages component-only reasoning. A model, platform, policy, assay, or infrastructure service acquires practical meaning only when its interfaces and use conditions are specified. The evidence chain should describe what information is available, what transformation occurs, which assumptions make that transformation credible, and who is authorized to act.

Lu et al. (2024) addresses "Uncertainty Quantification and Interpretability for Clinical Trial Approval Prediction." In this synthesis, that work is used as a source on system mechanism within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Mena et al. (2020) addresses "Uncertainty-Based Rejection Wrappers for Black-Box Classifiers." In this synthesis, that work is used as a source on representation choice within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Hu et al. (2023) addresses "Uncertainty in Natural Language Processing: Sources, Quantification, and Applications." In this synthesis, that work is used as a source on failure analysis within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

These studies also show why modular claims require interface tests. A component may perform well while the complete pathway fails because inputs drift, users interpret outputs differently, recovery semantics are ambiguous, or institutional incentives change. Design documentation should therefore include not only the intended pathway but also foreseeable deviations, degraded modes, and conditions under which the system should stop or defer.

4. Evaluation and Uncertainty

Evaluation should center on coverage-risk curves, calibration under shift, and error severity at the decision boundary. Average performance remains useful, but it should be accompanied by distributions, error categories, relevant subgroups or operating regimes, and uncertainty at the point where a decision changes. Comparators must isolate the value of the proposed addition rather than compare a mature intervention with an intentionally weak baseline.

Ehrlich-Sommer et al. (2025) addresses "ForestGPT and Beyond: A Trustworthy Domain-Specific Large Language Model Paving the Way to Forestry 5.0." In this synthesis, that work is used as a source on comparative evaluation within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Ren et al. (2023) addresses "Self-Evaluation Improves Selective Generation in Large Language Models." In this synthesis, that work is used as a source on uncertainty within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

External validation should introduce meaningful shifts in setting, time, population, workload, market structure, or environmental conditions. A nominally independent sample can still be too similar to test transfer. Conversely, a failed transfer is scientifically informative when the changed conditions are measured and the mechanism of failure is examined rather than hidden in an aggregate score.

5. Translation and Governance

Translation to high-consequence decision support with bounded model authority depends on more than technical readiness. Decision rights, documentation, maintenance, monitoring, appeal, and recovery are part of the intervention. The governance requirement is named escalation owners, traceable confidence signals, and post-release recalibration triggers. Each control should be connected to a named failure mode and should identify an owner, evidence source, review interval, and escalation path.

Geng et al. (2023) addresses "A Survey of Confidence Estimation and Calibration in Large Language Models." In this synthesis, that work is used as a source on implementation within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and prevents an adjacent result from being converted into a broader causal claim.

Tomani et al. (2024) addresses "Uncertainty-Based Abstention in LLMs Improves Safety and Reduces Hallucinations." In this synthesis, that work is used as a source on governance within calibrated abstention and escalation in foundation-model systems. The connection is deliberately bounded: the cited study informs the evidence map, but it does not by itself establish readiness for high-consequence decision support with bounded model authority. That distinction keeps the review close to the published record and

prevents an adjacent result from being converted into a broader causal claim.

A credible implementation plan also defines sunset and revalidation conditions. New data, software updates, institutional restructuring, population change, or shifts in incentives can invalidate previously reasonable assumptions. Monitoring should therefore test the validity of the evidence chain, not merely whether a service remains online or a headline metric remains stable.

6. Research Agenda

Future studies should preregister or otherwise predefine the intended decision, principal outcomes, exclusions, and analysis plan when feasible. They should report missingness, sensitivity analyses, negative or null findings, and the cost of errors to differently situated stakeholders. Reproducible artifacts should include sufficient metadata to reconstruct data provenance, model or analytical versions, parameter choices, and the sequence from evidence to conclusion.

Cross-disciplinary work needs a shared object of explanation rather than a collection of adjacent vocabularies. For calibrated abstention and escalation in foundation-model systems, that shared object is the complete decision pathway. Technical, clinical, social, environmental, or economic expertise should each change the design or interpretation of the study. When evidence remains incomplete, the useful output is a bounded hypothesis, a transparent uncertainty statement, or a request for additional measurement.

7. Conclusion

The literature supports a disciplined approach to calibrated abstention and escalation in foundation-model systems: define the decision, preserve system boundaries, test consequential failure modes, connect uncertainty to action, and assign accountable control. The strongest conclusion is not that one method is universally superior. It is that credible use in high-consequence decision support with bounded model authority requires an auditable link from source evidence to design choice, evaluation result, and governed decision.

References

- Ehrlich-Sommer, F., Eberhard, B., & Holzinger, A. (2025). ForestGPT and Beyond: A Trustworthy Domain-Specific Large Language Model Paving the Way to Forestry 5.0. *Electronics*, 14(18), 3583. <https://doi.org/10.3390/electronics14183583>
- Geng, J., Cai, F., Wang, Y., Koeppl, H., Nakov, P., & Gurevych, I. (2023). A Survey of Confidence Estimation and Calibration in Large Language Models. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2311.08298>
- Hu, M., Zhang, Z., Zhao, S., Huang, M., & Wu, B. (2023). Uncertainty in Natural Language Processing: Sources, Quantification, and Applications. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2306.04459>
- Lin, Z., Trivedi, S., & Sun, J. (2023). Generating with Confidence: Uncertainty Quantification for Black-box Large Language Models. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2305.19187>
- Lu, Y., Chen, T., Hao, N., Rechem, C. V., Chen, J., & Fu, T. (2024). Uncertainty Quantification and Interpretability for Clinical Trial Approval Prediction. *Health Data Science*, 4, 0126. <https://doi.org/10.34133/hds.0126>
- Mena, J., Pujol, O., & Vitria, J. (2020). Uncertainty-Based Rejection Wrappers for Black-Box Classifiers. *IEEE Access*, 8, 101721-101746. <https://doi.org/10.1109/access.2020.2996495>
- Ren, J., Zhao, Y., Vu, T., Liu, P. J., & Lakshminarayanan, B. (2023). Self-Evaluation Improves Selective Generation in Large Language Models. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2312.09300>
- Shorinwa, O., Mei, Z., Lidard, J., Ren, A. Z., & Majumdar, A. (2025). A Survey on Uncertainty Quantification of Large Language Models: Taxonomy, Open Research Challenges, and Future Directions. *ACM Computing Surveys*, 58(3), 1-38. <https://doi.org/10.1145/3744238>

- Tomani, C., Chaudhuri, K., Evtimov, I., Cremers, D., & Ibrahim, M. (2024). Uncertainty-Based Abstention in LLMs Improves Safety and Reduces Hallucinations. arXiv (Cornell University). <https://doi.org/10.48550/arxiv.2404.10960>
- Wen, B., Yao, J., Feng, S., Xu, C., Tsvetkov, Y., Howe, B., & Wang, L. L. (2025). Know Your Limits: A Survey of Abstention in Large Language Models. *Transactions of the Association for Computational Linguistics*, 13, 529-556. https://doi.org/10.1162/tacl_a_00754