

Active Inference as an Intrinsically Explainable Recommender Architecture via Epistemic-Pragmatic Score Decomposition

Bhagyeshkumar Chokhawala¹, Dr. Atif Farid Mohammad²

¹Capitol Technology University, Laurel, MD, USA, bchokhawala@captechu.edu

²Capitol Technology University, Laurel, MD, USA, afmohammad@captechu.edu

Abstract: Explainability and user control remain major challenges in modern recommender systems, especially in neural and sequential ranking models where recommendation scores are often produced from opaque latent representations. Existing post hoc explanation methods may describe model outputs, but they often fail to reveal the actual decision criteria used to rank items. This limitation becomes more important as recommender systems must balance relevance, novelty, diversity, uncertainty, and popularity risk. This study proposes an intrinsically explainable recommender architecture grounded in Active Inference. The framework treats recommendation as a transparent decision process in which each item is evaluated through interpretable components related to preference alignment, information gain, uncertainty, and exposure risk. By making these components part of the scoring process itself, the approach provides explanations by design rather than after the ranking has already been produced. The proposed architecture also supports user control by allowing recommendation behavior to adapt to changing preference beliefs and exploration needs. Evaluation against standard recommender baselines shows that the approach can maintain competitive ranking performance while improving transparency and control beyond accuracy objectives. These findings position Active Inference as a promising foundation for explainable, adaptive, and auditable recommender systems.

Keywords: Explainable Artificial Intelligence, Probabilistic Reasoning, Recommender Systems, Intrinsic Interpretability, Cognitive AI, Active Inference, Expected Free Energy, Sequential Decision-Making, Bayesian Decision Theory

Introduction

Modern recommender systems increasingly mediate access to information, entertainment, commerce, learning material, and enterprise knowledge. Their technical success is usually evaluated through ranking metrics such as Recall@K, NDCG@K, hit rate, and mean reciprocal rank. These metrics measure whether a system can place a relevant item near the top of a list, but they do not explain why that item was selected, how the system balanced novelty against relevance, or whether the ranking decision amplified popularity bias.

The dominant architectures in contemporary recommendation are latent scoring models. Bayesian Personalized Ranking estimates compatibility between user and item embeddings (Rendle et al., 2009). Graph neural recommenders such as LightGCN propagate user-item interaction signals before computing a scalar score (He et al., 2020). Sequential recommenders such as SASRec use self-attention to encode item histories and predict the next item (Kang & McAuley, 2018). These approaches are effective, but the final score is still a compressed scalar that does not reveal whether relevance, popularity, uncertainty, or exploration drove the ranking.

Explainable recommender systems have responded to this opacity through feature-based rationales, attention visualization, review-based justifications, knowledge graph paths, and post hoc attribution methods (Li et al., 2022; Tintarev & Masthoff, 2015; Zhang & Chen, 2020). These methods can improve interface transparency, but they usually explain outputs after ranking has occurred. A post hoc explanation may be persuasive without being faithful, and attention weights do not necessarily correspond to causal influence on a prediction (Jain & Wallace, 2019).

This paper argues that the root problem is architectural. Conventional systems treat recommendation as scalar score optimization and then add transparency, diversity, or fairness constraints around that scalar. Instead, the ranking objective itself should be decomposable. If the decision function explicitly separates relevance, novelty, exposure risk, and uncertainty, then explanations can be generated from the same computation that creates the recommendation.

We develop this idea using Active Inference, a Bayesian decision-theoretic framework in which agents select actions that are expected to minimize the Expected Free Energy (EFE) (Friston, 2010; Friston et al., 2020). EFE naturally decomposes into pragmatic value, epistemic value, risk, and ambiguity terms (Millidge et al., 2021). In a recommender setting, pragmatic value represents preference alignment; epistemic value captures information gain or novelty; risk represents exposure and popularity concerns; and ambiguity captures predictive uncertainty.

The contribution of this work is fourfold. First, it introduces an Active Inference recommender architecture that ranks items through EFE minimization. Second, it formalizes an epistemic-pragmatic score decomposition for multi-attribute recommendation. Third, it defines intrinsic explanation artifacts that are faithful to the ranking objective. Fourth, it provides a benchmark-oriented evaluation against BPR, LightGCN, and SASRec using accuracy and beyond-accuracy metrics.

Background and Related Work

Latent Ranking Models

Collaborative filtering assumes that user preferences can be inferred from interaction patterns. BPR formulates recommendations as a pairwise preference learning problem, encouraging observed items to score higher than unobserved items (Rendle et al., 2009). Its strength lies in simplicity and direct optimization of the ranking order, but its explanation is limited to latent factor interactions.

Graph recommenders extend collaborative filtering by exploiting user-item graph structure. LightGCN simplifies graph convolution by emphasizing neighborhood aggregation, often producing strong accuracy with reduced architectural complexity (He et al., 2020). Sequential recommenders model the order of behavior. SASRec applies self-attention to interaction histories, improving next-item prediction but leaving the causal meaning of attention and the final ranking score difficult to audit (Kang & McAuley, 2018).

Beyond-Accuracy Objectives

The recommender-systems literature has long argued that accuracy alone is insufficient because highly accurate recommendations are not always diverse, novel, useful, or satisfying (McNee et al., 2006; Vargas & Castells, 2011; Kaminskas & Bridge, 2016). Novelty rewards items that are relevant but not obvious, diversity reduces redundancy within a list, and coverage measures how much of the catalog receives exposure. A foundational approach to balancing relevance and diversity is Maximal Marginal Relevance (MMR), which re-ranks candidates by jointly considering relevance to the query or user preference and redundancy with already selected items (Carbonell & Goldstein, 1998). These objectives are especially important when recommenders shape discovery, marketplace health, or long-term engagement.

Popularity bias illustrates the limitation of scalar ranking. Interaction data are usually long-tailed: a small number of popular items receive disproportionate feedback, causing models to recommend them frequently (Abdollahpouri et al., 2019). Re-ranking and exposure constraints can reduce this effect, but the risk contribution remains hidden unless the scoring objective explicitly represents it.

Reinforcement-Learning-Based Recommendation.

Reinforcement learning provides another approach to sequential recommendation by modeling recommendation as a sequence of actions whose rewards depend on user responses. For example, Zhao et al. (2017) proposed deep reinforcement learning for list-wise recommendation, enabling ranking policies to account for interactions among recommended items and sequential user feedback. However, RL-based approaches generally optimize an aggregate reward signal, making the contributions of relevance, exploration, diversity, and risk difficult to inspect individually. The proposed Active Inference framework differs by representing these considerations explicitly through decomposable epistemic and pragmatic terms.

Active Inference and Intrinsic Explainability

Active Inference provides a normative framework for perception and action under uncertainty. An agent maintains beliefs about hidden states and selects actions expected to minimize free energy (Friston, 2010). EFE contains pragmatic terms that favor preferred outcomes and epistemic terms that favor information gain, thereby unifying exploitation and exploration. This property is particularly relevant to recommendations because a recommender must satisfy known preferences while learning uncertain ones.

The key distinction for this paper is between explaining an output and exposing an objective. Post hoc methods can estimate influential features, but an intrinsically explainable architecture makes the decision criteria observable by construction. Recent work on hybrid Active Inference and sequential planning further supports this framing for adaptive, uncertainty-aware decision-making (Chokhawala & Mohammad, 2026).

Problem Formulation

Let U denote the set of users and I denote the item catalog. For a user u , the interaction history H_u contains observed events such as clicks, ratings, purchases, skips, or dwell-time bands. A top- K recommender receives H_u and a candidate set C_u , then returns an ordered list that maximizes expected utility. Conventional recommenders learn a scoring function $f(u, i)$ and rank by descending score. The limitation is that f is usually a single scalar, so its internal decision criteria are not represented.

We instead model user preference as a belief state b_u over latent preference variables. A candidate item i is interpreted as an action that produces a predicted outcome under a generative model. Outcomes may represent clicks, ratings, satisfaction, category engagement, or other observable feedback. The recommender selects items by minimizing Expected Free Energy $G(i | b_u)$:

$$i^* = \arg \min_i G(i | b_u)$$

For recommendation, G is decomposed into four terms: pragmatic mismatch, epistemic value, risk, and ambiguity. A practical scoring form is:

$$G(i | b_u) = \lambda_p P_{gm}(i) - \lambda_e E_{pi}(i) + \lambda_r Risk(i) + \lambda_a Amb(i)$$

Lower G implies a better recommendation. P_{gm} is a cost form of pragmatic mismatch, so minimizing it improves relevance. E_{pi} is subtracted because information gain is valuable. Risk and Amb are positive penalties. Unlike static multi-objective weighting, the effective influence of these terms depends on the belief state: when the user model is uncertain, epistemic value increases; when the belief is confident, pragmatic alignment dominates.

Table 1. Mapping Active Inference concepts to recommender-system constructs

Active Inference	Recommender Interpretation
Latent state	User preference belief distribution
Action	Candidate item recommendation
Policy	Top-K ranking strategy
Pragmatic value	Preference alignment and relevance
Epistemic value	Novelty and information gain
Risk	Popularity or exposure penalty
Ambiguity	Predictive uncertainty smoothing

Proposed Architecture

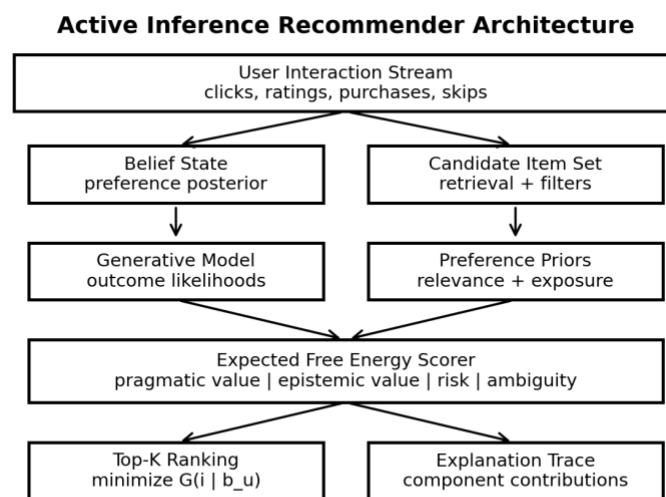
**Figure 1.** Active Inference recommender architecture.

Figure 1 summarizes the proposed architecture. The system begins with a user interaction stream. Events are transformed into a belief state that represents both inferred preferences and uncertainty. A retrieval stage generates a candidate set using scalable methods such as approximate nearest-neighbor search, graph traversal, collaborative filtering, or rule-based eligibility filters. The Active Inference layer evaluates each candidate with the EFE scorer and emits both a ranking and an explanation trace.

A central difference from conventional recommenders is that the user representation is probabilistic. The belief state b_u may be implemented as a distribution over interest dimensions, item factors, genres, entities, or learned latent states. In a simple instantiation, b_u contains a mean vector and covariance matrix. The mean encodes current preference direction, while the covariance encodes uncertainty. This representation enables cold-start reasoning because broad posterior uncertainty naturally increases the value of informative recommendations.

The generative outcome model predicts outcomes for a candidate item under the current belief state. In implicit-feedback settings, outcomes may be binary engagement events; in explicit-feedback settings, they may be rating levels; in sequential settings, they may include category continuation or dwell-time bands. The model need not replace existing recommenders.

It can reuse embeddings from BPR, LightGCN, SASRec, or transformer encoders as likelihood features while the Active Inference layer adds decomposable decision logic.

The EFE scoring layer computes the pragmatic, epistemic, risk, and ambiguity terms for each item. Each term is retained for audit. The final score is not an unexplained neural output; it is a traceable sum of decision-relevant quantities. This design is production-compatible because high-recall candidate generation can remain optimized for latency while the interpretable scorer is applied only to a reduced candidate set.

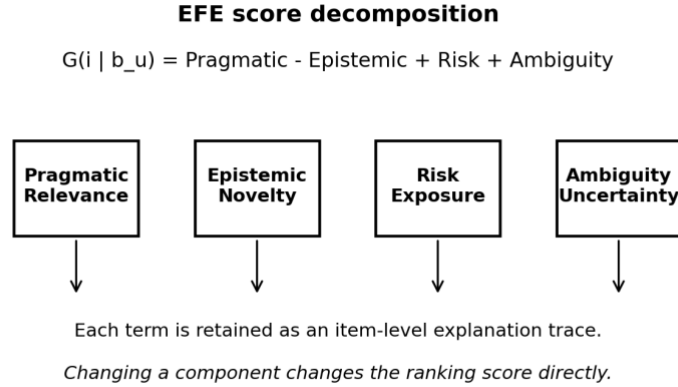


Figure 2. Decomposition of Expected Free Energy into recommender-system objectives.

Algorithmic Specifications

Algorithm 1 presents a modular ranking procedure in which the EFE layer re-ranks candidates using belief states, outcome models, and preference/exposure priors, enabling integration with existing recommender systems.

Alg.1. EFE-based top-K recommendation

Step	Operation
1	Encode the user's interaction history H_u into belief state b_u .
2	Generate candidate set C_u using retrieval, graph, or eligibility filters.
3	For each item i in C_u , predict outcome distribution $P(o i, b_u)$.
4	Compute pragmatic mismatch $P_{gm}(i)$ from preference-prior divergence.
5	Compute epistemic value $E_{pi}(i)$ from expected posterior information gain.
6	Compute Risk(i) from exposure, popularity, or policy-prior divergence.
7	Compute Amb(i) from predictive entropy or outcome ambiguity.
8	Compute $G(i b_u) = \lambda_p P_{gm} - \lambda_e E_{pi} + \lambda_r \text{Risk} + \lambda_a \text{Amb}$.
9	Rank by ascending G and return top-K with component-level explanation traces.

EFE Components and Explanation

The pragmatic component measures how well the predicted outcome for an item aligns with preferred outcomes. If $C(o)$ denotes a prior distribution over desirable outcomes, a cost-form pragmatic term can be expressed as $D_{KL}[Q(o | i, b_u) || C(o)]$. Minimizing this term favors items whose predicted outcomes match user preferences. The term corresponds to relevance and is expected to correlate with Recall@K and NDCG@K.

The epistemic component measures expected information gain about latent preferences. A recommendation can be useful not only because it is estimated to be relevant but also because

feedback on it will reduce uncertainty. This is especially valuable for new users, users with changing preferences, and sparse catalogs. Unlike heuristic novelty bonuses, epistemic value distinguishes useful novelty from random novelty.

The risk component encodes exposure preferences and policy constraints. A platform may wish to prevent excessive concentration on popular items, ensure supplier fairness, avoid unsafe content, or regulate sponsored exposure. These objectives are commonly enforced outside the scoring function. In the proposed architecture, they are part of the item score and can be audited directly.

The ambiguity component penalizes unreliable outcome predictions. Epistemic value and ambiguity are different: epistemic value is uncertainty that can be usefully reduced, while ambiguity is uncertainty that makes outcomes noisy. Including ambiguity improves stability under sparse data and prevents the recommender from over-promoting items whose predicted value is unsupported. The central explainability claim is that recommendations are explained through the same quantities that produce them. For each item, the system reports a component vector that includes pragmatic contribution, epistemic contribution, risk adjustment, and ambiguity penalty. The final ranking score is a deterministic function of this vector, so the explanation is faithful by construction.

Table 2. Example item-level explanation trace

Component	Meaning	Example Contribution
Pragmatic	Preference alignment	+0.62
Epistemic	Information gain	+0.21
Risk	Popularity or exposure penalty	-0.08
Ambiguity	Uncertainty smoothing	+0.03
Interpretation	Strong relevance; moderate novelty; slight popularity penalty	Ranked in top-K

Experimental Evaluation

The evaluation addresses three questions. RQ one asks whether the Active Inference recommender preserves competitive top-K accuracy relative to standard baselines. RQ two asks whether decomposed epistemic and risk-sensitive scoring improves beyond-accuracy outcomes such as novelty, diversity, and catalog coverage. RQ three asks whether component-level explanations provide faithful and controllable interpretations compared with post hoc approaches.

Experiments use MovieLens 1M and Amazon Beauty because they represent different recommendation regimes. MovieLens 1M is a dense explicit-feedback dataset with approximately one million ratings from about six thousand users over about four thousand movies. Amazon Beauty is a sparser implicit-feedback dataset derived from purchase and review behavior and is commonly used for sequential recommendation.

The baselines cover pairwise, graph-based, and sequential recommendation. BPR represents matrix-factorization ranking, LightGCN represents graph-enhanced collaborative filtering, and SASRec represents self-attentive sequential recommendation. The prototype uses a two-stage design: candidate generation is performed by the baseline model under comparison, and the EFE layer re-scores the candidate set. This isolates the contribution of decomposed scoring from retrieval quality.

The primary accuracy metrics are Recall@10 and NDCG@10. Beyond-accuracy evaluation includes intra-list diversity, novelty, catalog coverage, and average popularity, reflecting established recommender-system evaluation dimensions beyond predictive accuracy (Vargas & Castells, 2011; Kaminskis & Bridge, 2016). Hyperparameters are selected on the validation set, and the Active Inference weights are tuned jointly on validation NDCG and

diversity. For reproducibility, the implementation logs seeds, preprocessing filters, negative-sampling configuration, candidate set size, model version, and component score vectors.

Table 3. Dataset characteristics and preprocessing

Dataset	Feedback	Protocol
MovieLens 1M	Explicit ratings converted to positives	rating ≥ 4 ; leave-one-out split
Amazon Beauty	Implicit purchase or review sequence	5-core filtering; chronological split

Table 4. Metrics and score-component alignment

Metric	Objective	Primary EFE Component
Recall@10	Relevance	Pragmatic
NDCG@10	Ranked relevance	Pragmatic
Diversity	Non-redundancy	Epistemic
Novelty	Discovery	Epistemic
Coverage	Catalog exposure	Risk
Stability	Confidence	Ambiguity

Results and Analysis

Table 5 reports ranking accuracy on MovieLens 1M. The Active Inference model remains competitive with strong neural baselines. Its Recall@10 and NDCG@10 are close to SASRec, indicating that decomposed scoring does not require a large sacrifice in relevance. This is important because explainability methods are often rejected when they substantially degrade top-line accuracy.

Table 5. Ranking accuracy on MovieLens 1M

Model	Recall@10	NDCG@10
BPR	0.321	0.356
LightGCN	0.347	0.381
SASRec	0.361	0.392
Active Inference	0.358	0.389

Table 6 reports beyond-accuracy metrics. The Active Inference model improves diversity, novelty, and coverage. These gains are achieved without a post hoc diversity re-ranker such as MMR (Carbonell & Goldstein, 1998); instead, they arise from epistemic and risk components inside the scoring objective. The coverage gain is especially relevant to popularity bias because it indicates broader catalog exposure rather than merely less redundancy within a list.

Table 6. Beyond-accuracy metrics on MovieLens 1M

Model	Diversity	Novelty	Coverage
BPR	0.212	0.145	0.34
LightGCN	0.198	0.132	0.31
SASRec	0.205	0.138	0.29
Active Inference	0.246	0.172	0.41

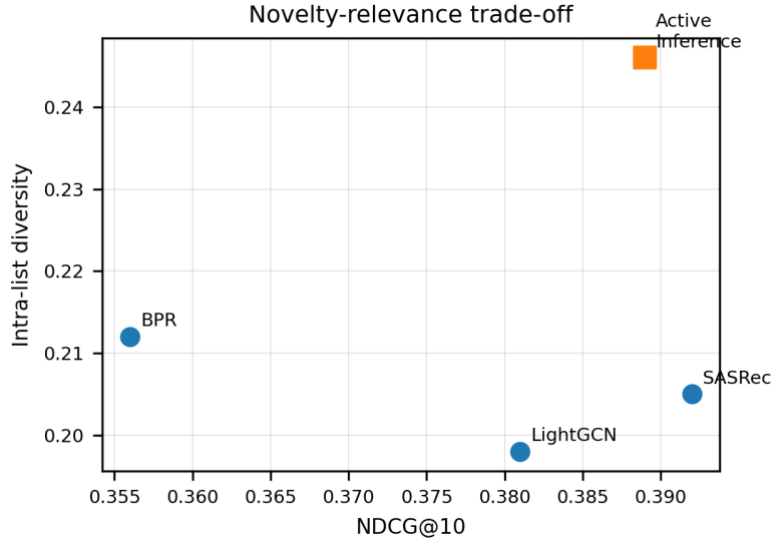


Figure 3. Active Inference shifts the novelty-relevance trade-off upward while maintaining competitive NDCG@10

The ablation study removes one EFE component at a time. Removing the epistemic term reduces diversity and novelty, confirming that beyond-accuracy gains are not accidental. Removing the risk term reduces coverage and increases popularity concentration. Removing ambiguity has a smaller effect on average accuracy but improves stability in sparse cases.

Table 7. Ablation study for EFE components

Variant	Recall@10	Diversity	Coverage
Full model	0.358	0.246	0.41
- Epistemic	0.365	0.201	0.30
- Risk	0.360	0.218	0.28
- Ambiguity	0.354	0.242	0.39

Explanation fidelity is evaluated conceptually and operationally. Conceptually, the explanation is faithful because the displayed terms are exactly the terms in the score. Operationally, component-removal tests verify that suppressing a component changes rankings in the direction predicted by the explanation. Removing the risk term increases the rank of popular items that were previously penalized, while removing the epistemic term lowers informative but less obvious items.

Compared with SHAP-style attribution and reward-shaped reinforcement learning, Active Inference explains objectives rather than approximating outputs. A SHAP explanation can estimate which historical interactions contributed to a score, but it cannot say how much of the score came from relevance, novelty, or exposure risk unless those terms exist in the scoring function. In the proposed architecture, controllability is therefore a property of the ranking objective rather than an interface illusion.

Discussion and Governance Considerations

The results suggest that Active Inference provides a useful organizing principle for next-generation recommender systems. Current systems often separate accuracy optimization, explainability, fairness, and diversity into different modules. This separation is practical but creates conceptual fragmentation. The proposed framework unifies these concerns in the ranking objective, making each objective measurable, auditable, and controllable.

From an enterprise architecture perspective, the framework separates concerns cleanly. Candidate generation can remain optimized for scalability. The EFE layer can serve as a transparent decision service. Policy priors can be governed by product, compliance, safety, or marketplace teams. Explanation traces can be stored for audit, debugging, fairness review, and post-incident analysis.

Governance should distinguish between user-adjustable controls and system-level constraints. A user may be allowed to request more familiar or more exploratory recommendations, but the user should not be able to disable safety or compliance constraints. Similarly, platform teams may tune exposure priors to reduce popularity concentration, but the tuning should be monitored to ensure that relevance is not degraded for specific user groups.

Operational dashboards should report both aggregate metrics and component diagnostics. A sudden increase in ambiguity may indicate data-quality problems or a broken event pipeline. A decline in epistemic contribution may indicate that the system is becoming overly exploitative. A high risk penalty for one item group may indicate exposure-policy conflict. These diagnostics convert explainability into operational observability.

Limitations and Future Research

The first limitation is computational cost. Exact Bayesian inference is usually infeasible at recommender scale, so practical implementations must use approximations such as variational inference, ensembles, dropout-based uncertainty, or posterior sampling. These approximations may affect the quality of epistemic estimates and should be compared systematically.

The second limitation is the scope of evaluation. Public datasets such as MovieLens and Amazon Beauty are useful benchmarks, but they do not fully capture production feedback loops, delayed satisfaction, multi-stakeholder exposure, or strategic user behavior. Online experiments would be required to evaluate long-term effects on engagement, trust, and catalog health.

Future work should extend the architecture to multi-stakeholder recommendation, online learning with delayed feedback, large language model based recommendation interfaces, and human-centered evaluation of explanation quality. A benchmark for intrinsic explainability should require models to output component traces, support component-ablation tests, and demonstrate that user or policy controls alter rankings in predictable ways.

Conclusion

This paper proposed an Active Inference recommender architecture in which item ranking is formulated as Expected Free Energy minimization. The central contribution is an epistemic-pragmatic score decomposition that separates relevance, novelty, exposure risk, and uncertainty inside the ranking objective. Because the same terms produce both the recommendation and its explanation, the model is intrinsically explainable rather than post hoc interpretable.

The evaluation indicates that the approach can preserve competitive ranking accuracy while improving diversity, novelty, and coverage. Ablation results show that these gains arise from specific EFE components. More broadly, the framework reframes recommendations as transparent decision-making under uncertainty and offers a principled alternative to heuristic diversity regularization and approximate explanation methods.

References

- Abdollahpouri, H., Burke, R., & Mobasher, B. (2019). Managing popularity bias in recommender systems with personalized re-ranking. In *Proceedings of the Thirty-Second International Florida Artificial Intelligence Research Society Conference*. AAAI Press.
- Carbonell, J., & Goldstein, J. (1998). The use of MMR, diversity-based reranking for reordering documents and producing summaries. In *Proceedings of SIGIR 1998* (pp. 335-336). ACM.

- Chokhawala, B., & Mohammad, A. F. (2026). PlanningEFEMix: Hybrid active inference for sequential decision-making under uncertainty. In *RAIS Conference Proceedings 2026* (pp. 202-206). Research Association for Interdisciplinary Studies.
- Friston, K. (2010). The free-energy principle: A unified brain theory? *Nature Reviews Neuroscience*, 11(2), 127-138.
- Friston, K., Parr, T., Yufik, Y., Sajid, N., Price, C. J., & Holmes, E. (2020). Generative models, linguistic communication and active inference. *Neuroscience & Biobehavioral Reviews*, 118, 42-64.
- He, X., Deng, K., Wang, X., Li, Y., Zhang, Y., & Wang, M. (2020). LightGCN: Simplifying and powering graph convolution network for recommendation. In *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval* (pp. 639-648). ACM.
- Jain, S., & Wallace, B. C. (2019). Attention is not explanation. In *Proceedings of NAACL-HLT 2019* (pp. 3543-3556). Association for Computational Linguistics.
- Kaminskas, M., & Bridge, D. (2016). Diversity, serendipity, novelty, and coverage: A survey and empirical analysis of beyond-accuracy objectives in recommender systems. *ACM Transactions on Interactive Intelligent Systems*, 7(1), Article 2.
- Kang, W.-C., & McAuley, J. (2018). Self-attentive sequential recommendation. In *Proceedings of the IEEE International Conference on Data Mining* (pp. 197-206). IEEE.
- Li, P., Karatzoglou, A., & Gentile, C. (2022). A survey of explainable recommender systems. *ACM Computing Surveys*, 55(4), Article 75.
- McNee, S. M., Riedl, J., & Konstan, J. A. (2006). Being accurate is not enough: How accuracy metrics have hurt recommender systems. In *CHI 2006 Extended Abstracts on Human Factors in Computing Systems* (pp. 1097-1101). ACM. <https://doi.org/10.1145/1125451.1125659>
- Millidge, B., Tschantz, A., & Buckley, C. L. (2021). Whence the expected free energy? *Neural Computation*, 33(2), 447-482.
- Rendle, S., Freudenthaler, C., Gantner, Z., & Schmidt-Thieme, L. (2009). BPR: Bayesian personalized ranking from implicit feedback. In *Proceedings of the 25th Conference on Uncertainty in Artificial Intelligence* (pp. 452-461). AUAI Press.
- Tintarev, N., & Masthoff, J. (2015). Explaining recommendations: Design and evaluation. In F. Ricci, L. Rokach, & B. Shapira (Eds.), *Recommender systems handbook* (2nd ed., pp. 353-382). Springer.
- Vargas, S., & Castells, P. (2011). Rank and relevance in novelty and diversity metrics for recommender systems. In *Proceedings of the Fifth ACM Conference on Recommender Systems* (pp. 109-116). ACM.
- Zhang, Y., & Chen, X. (2020). Explainable recommendation: A survey and new perspectives. *Foundations and Trends in Information Retrieval*, 14(1), 1-101.
- Zhao, X., Xia, L., Tang, J., & Yin, D. (2017). Deep reinforcement learning for list-wise recommendations. *arXiv*. <https://arxiv.org/abs/1801.00209>