

Active Inference Mapping for an Event-Spine World Model

Core correspondence

Your three-layer design maps to active inference cleanly if the event spine is treated as the realized history of observations, actions, and outcomes; the graph layer is treated as a materialized factorization of current hidden causes; and the belief layer is treated as the inferential layer that maintains posterior beliefs, precisions, and model-comparison scores over those causes. In the discrete active-inference formulation, agents infer hidden states and events from observations, maintain beliefs about how states evolve over time, and act on the basis of those beliefs rather than raw observations. On that reading, a planner that does **not** read raw facts but instead reads a settled summary is compatible with the literature. ¹

The main mismatch is architectural rather than conceptual. In canonical active inference, state estimation and policy selection are not two unrelated modules: current-state inference minimizes variational free energy, while policy selection minimizes expected free energy or generalized free energy. A separate planner remains viable as an engineering decision, but then the BeliefView should carry policy-relevant structure, not only settled state labels. At minimum, it should include posterior beliefs over anchors, uncertainty or precision on those beliefs, and whether the next lowest-free-energy move is to **act** or to **observe**. ²

A concise translation of your loop into active-inference terms is: observe the world; register realized outcomes in the spine; update posterior beliefs over hidden causes and anchors; evaluate policies over future outcomes; emit a planner-facing belief or policy view; act; and then write the realized consequences back as new outcomes. That translation is especially strong because the generalized-free-energy formulation treats past outcomes as observations and future outcomes as hidden states that will only be disclosed by time. ³

Markov blanket placement

The answer to **(a)** is: the Markov blanket does **not** sit at the graph boundary or the BeliefView boundary. In the FEP literature, the blanket is the **system–environment interface**. It is the set of blanket states that mediate conditional independence between internal and external states, and it is partitioned into **sensory** states and **active** states. Sensory states are influenced by the environment; active states influence the environment. That is a statistical-interface concept, not a storage-layer concept. ⁴

Applied to your stack, the closest analogue is the **edge around spine ingress and egress**. Incoming observations that are serialized into the spine are the closest analogue of sensory blanket states; outbound task executions and their world-facing effects are the closest analogue of active blanket states. The spine, graph, and belief layers themselves are better interpreted as **internal** states or internal-state machinery.

Once an observation has been normalized into an immutable semantic fact, it has already crossed the blanket and been internalized. ⁵

If forced to choose only one of your proposed boundaries, the least-wrong answer is **the spine boundary**, but only in the narrow sense of the ingest and emit adapters at the edge of the spine. The fact log itself is not the blanket. The graph layer is a representation of inferred hidden causes. The BeliefView is a planner-facing posterior summary. Neither one mediates the world-facing conditional-independence relation that defines a Markov blanket in the literature. ⁴

Belief revision as free energy minimization

The answer to **(b)** is that “compute a posterior” and “minimize variational free energy” are not alternatives of the same type. The posterior is the **target belief state**. Variational free energy is a **computational objective** used when exact Bayesian inference is not tractable. The active-inference tutorial literature is explicit that variational free energy is used for approximate Bayesian inference, and that minimizing it minimizes the divergence between the approximate posterior and the true posterior. Therefore, if your comparator can compute an exact posterior over anchor credibility in a tractable model, that is already compatible with the framework. If not, minimizing variational free energy is the standard route. ⁶

Under conflicting evidence, the active-inference move is to interpret the comparator as carrying out approximate inference and model comparison over competing hidden causes, anchors, or perspectives. The output of that process should still be posterior-like: belief mass, precision, and conflict diagnostics. The free-energy scalar is the optimization criterion, not the planner-facing semantic artifact. In practical terms: use **variational free energy** inside the comparator, and let the **BeliefView** expose posterior marginals and uncertainty. ⁷

The place where plain posterior computation becomes insufficient is **policy selection about the future**. The generalized-free-energy paper shows that the same form of posterior policy updating can be obtained under both expected-free-energy and generalized-free-energy formulations, while differing in how future observations are represented and scored. If your comparator is only revising current anchor credibility, a standard posterior or approximate posterior is sufficient. If the comparator is also choosing whether to observe, wait, or intervene, then it should score policies using **expected free energy** or, if you want a single objective spanning present and future, **generalized free energy**. ⁸

A compact implementation rule follows from this. Use **VFE** for present-state inference and anchor settlement. Use **EFE/GFE** for selecting policies that affect future observations, including information-seeking actions. That separation matches the literature more closely than replacing the posterior with a free-energy number. ⁷

Generative model for deep temporal uncertainty

The answer to **(c)** is that the FEP literature does **not** recommend a flat current-state model when future effects are unobservable. It recommends a **deep temporal generative model**. In the classic deep-temporal review, higher-level states entail sequences of transitions at lower levels, so evidence can accumulate across nested timescales. In the generalized-free-energy formulation, future outcomes are treated as latent or hidden states that will only be revealed by the passage of time. In sophisticated-inference work, planning is

extended to recursive beliefs about counterfactual consequences, producing a deep tree search over future actions and outcomes. ⁹

For software design, that implies at least five structural requirements in the generative model behind your belief layer. First, explicit hidden-state factors for the world conditions that anchors summarize. Second, explicit observation models for each event modality. Third, policy-conditioned transition models over time. Fourth, prior preferences over future observations or outcomes. Fifth, temporal abstraction: higher levels should summarize slower narrative or task structure, while lower levels track faster event dynamics. That is the recurring structure across discrete POMDP formulations, deep temporal models, hierarchical navigation models, and hybrid active-inference controllers. ¹⁰

Recent engineering work pushes this further in three directions. One direction uses **recursive tree search** or branching-time formulations for deep counterfactual planning. A second uses **dynamic programming over expected free energy** to reduce computational cost in finite-horizon discrete settings. A third uses **paths or trajectories as latent variables**, so the model represents not only states but lawful multi-step evolutions and events with temporal depth. Those are all direct responses to the problem you stated: the present action has future effects that are not yet observable. ¹¹

The most direct implication for your architecture is that the graph layer should remain a current-state materialization only if the belief layer reasons over **trajectories**, not just anchors. A graph anchor such as “subject points to target under perspective” can be the present-time slice of a deeper generative structure, but the comparator should score how that anchor evolves under policies and observations over time. The recent hybrid-model literature is explicit on this: the agent maintains potential trajectories, and a high-level discrete model uses them to infer world state and plan composite action in dynamic environments. ¹²

For a system like yours, a stronger active-inference alignment would be:

- event spine as realized past outcomes and actions;
- graph anchors as current materialized hypotheses;
- belief layer as posterior inference over hidden states **and** policies;
- BeliefView as planner-facing posterior state plus uncertainty and observation intents;
- task planner as executor or compiler of already ranked policies, not a separate reward optimizer over raw facts. ¹³

Epistemic value and observation requests

The answer to **(d)** is **yes**, with a qualification: unresolved beliefs should generate observation requests **when** the expected information gain of observing is greater than the cost and delay of observing. In the active-inference literature, epistemic value is the mutual information between hidden states and observations, and minimizing expected free energy directs the agent toward policies that reduce ambiguity and increase information gain. In direct terms, “needs_observation” is a software-friendly encoding of **positive epistemic value**. ¹⁴

This does **not** mean that every weakly supported belief should trigger a query. EFE combines pragmatic and epistemic terms. Therefore, an observation request should be emitted only when the unresolved belief is decision-relevant, the system has an available channel for disambiguation, and the expected reduction in uncertainty matters enough relative to action cost, latency, or opportunity cost. The constrained-Bethe-free-

energy work is especially relevant here because it models future observation explicitly by encoding the assumption that the agent **will** make observations in the future. That is a strong formal precedent for treating observation requests as first-class actions. ¹⁵

A direct design consequence follows. The belief layer should not only emit statuses like `settled`, `contested`, or `needs_observation`; it should emit a **candidate observation policy** with at least: target belief, candidate observation channel, expected ambiguity reduction, expected pragmatic cost, and expiry horizon. In active-inference terms, the observation request is not an exception path. It is a policy selected because sampling is predicted to lower expected free energy more than acting immediately under ambiguity. ¹⁶

That mapping is especially strong for your architecture because the planner already reads only the BeliefView. The belief layer can therefore decide whether the next planner-facing object is a settled state summary or an epistemic action request. That keeps the planner away from raw facts **and** makes unresolved beliefs actionable in a principled way. ¹⁷

Software-oriented literature from 2022 to 2025

Direct work on “agentic software systems” under FEP remains smaller than the neuroscience and robotics literature. The strongest matches to your requirements are the papers that make discrete state spaces, event-conditioned observation, reactive execution, or multi-agent coordination explicit.

Closest technical fits for your architecture

- **Active Inference and Epistemic Value in Graphical Models** (2022). This is one of the closest matches to your belief layer: free-form graphical models, message passing, and an explicit treatment of epistemic behavior through constrained Bethe free energy. It is especially relevant to `needs_observation`, because it encodes future observation as an explicit assumption of the planning objective. ¹⁸
- **On efficient computation in active inference** (2024). This paper introduces Dynamic Programming Expected Free Energy, using Bellman-style reverse recursion to reduce planning complexity in grid-world settings. It is a direct candidate for replacing an exhaustive comparator or planner over long discrete horizons. ¹⁹
- **On Predictive Planning and Counterfactual Learning in Active Inference** (2024). This paper compares planning-based and experience-based active-inference schemes and adds a mixed model in a grid-world scenario. It is useful if your architecture needs both explicit counterfactual planning and replay-based adaptation. ²⁰
- **From pixels to planning: scale-free active inference** (2025). This is the strongest recent match for discrete state spaces and event abstractions. It generalizes POMDPs by adding **paths** as latent variables, builds event-based representations, and applies the framework to Atari-like games. For an event-spine system, this is a strong template for representing temporally extended events rather than only instantaneous states. ²¹

- **Deep Active Inference Agents for Delayed and Long-Horizon Environments** (2025). This work adds multi-step latent transitions, an integrated policy network, and long-horizon planning in a delayed industrial environment. It is one of the clearest “world-model style” active-inference papers in the period you requested. ²²
- **Active Inference and Behavior Trees for Reactive Action Planning and Execution in Robotics** (2023). This hybrid of behavior trees and active inference is relevant if your planner compiles tasks and then executes them reactively. It handles partial observability, online adaptation, and runtime contingency handling while preserving hierarchical execution structure. ²³
- **Spatial and Temporal Hierarchy for Autonomous Navigation Using Active Inference in Minigrid Environment** (2024). This paper is valuable for discrete hierarchical planning: it uses a MiniGrid setting, combines curiosity-driven exploration with goal-oriented navigation, and explicitly plans across context, place, and motion levels. ²⁴
- **Integrating Cognitive Map Learning and Active Inference for Planning in Ambiguous Environments** (2023/2024). This is a strong fit when anchors are ambiguous or aliased, because it combines a learned cognitive-map structure with active-inference planning under uncertain location information. ²⁵

Papers emphasizing multi-agent coordination and belief sharing

- **Interactive Inference: A Multi-Agent Model of Cooperative Joint Actions** (journal version 2024; preprint 2022). This is a direct coordination paper: agents maintain probabilistic beliefs about joint goals, update those beliefs from other agents’ movements, and choose actions that make their own intentions legible. It is highly relevant if your anchors include beliefs about others’ intentions or commitments. ²⁶
- **Federated inference and belief sharing** (2024). This is the clearest recent paper on distributed belief sharing under free-energy minimization. It addresses how multiple agents with a shared world model communicate beliefs about a common world, which is directly relevant to multi-agent spine synchronization or cross-agent anchor credibility. ²⁷
- **Factorised Active Inference for Strategic Multi-Agent Interactions** (AAMAS 2025). Each agent maintains explicit beliefs about the internal states of other agents and uses them for strategic planning in non-stationary games. This is the strongest 2025 fit for multi-agent anchor scoring and perspective-indexed beliefs. ²⁸

Adjacent agentic-system papers

- **Active Inference for Self-Organizing Multi-LLM Systems** (2024/2025 preprint). This places an active-inference cognitive layer above an LLM-based agent, with state factors for prompt, search, and information, and multiple observation modalities for quality metrics. It is a direct attempt to port active inference into software agents, though it is a preprint and the empirical scope is narrower than the robotics and planning papers above. ²⁹

- **Agentic rulebooks using active inference** (2025). This is directly about autonomous AI agents resolving conflicting norms under uncertainty. It is less about discrete observation dynamics than the planning papers, but it is relevant if your planner must compile tasks subject to policy constraints or rule layers instead of scalar rewards. ³⁰
- **Active Digital Twins via Active Inference** (2025 preprint; journal version later). This is relevant if your event spine functions like an operational digital twin. It formulates the digital twin as a POMDP-style active-inference agent that updates beliefs online and chooses actions balancing pragmatics and epistemics. It is more engineering-oriented than software-agent-oriented, but it is a useful adjacent reference. ³¹

Taken together, the most transferable set for your design is: graphical-model and message-passing work for the comparator; DPEFE, sophisticated inference, and long-horizon active-inference agents for planning under delayed effects; federated, interactive, and factorized inference for multi-agent belief sharing; and scale-free or hybrid models for event-like abstractions that sit above raw facts and below planning. ³²

¹ ² ⁶ ⁷ ¹⁰ ¹⁴ ¹⁶ <https://www.repository.cam.ac.uk/bitstreams/d3022cd7-4ea2-4a50-a501-f515015c97fb/download>

<https://www.repository.cam.ac.uk/bitstreams/d3022cd7-4ea2-4a50-a501-f515015c97fb/download>

³ ⁸ ¹³ ¹⁷ <https://link.springer.com/article/10.1007/s00422-019-00805-w>

<https://link.springer.com/article/10.1007/s00422-019-00805-w>

⁴ ⁵ <https://arxiv.org/abs/2106.13830>

<https://arxiv.org/abs/2106.13830>

⁹ <https://www.sciencedirect.com/science/article/pii/S0149763418302525>

<https://www.sciencedirect.com/science/article/pii/S0149763418302525>

¹¹ <https://arxiv.org/abs/2111.11107>

<https://arxiv.org/abs/2111.11107>

¹² <https://arxiv.org/html/2402.10088v4>

<https://arxiv.org/html/2402.10088v4>

¹⁵ ¹⁸ ³² <https://arxiv.org/abs/2109.00541>

<https://arxiv.org/abs/2109.00541>

¹⁹ <https://www.sciencedirect.com/science/article/pii/S0957417424011813>

<https://www.sciencedirect.com/science/article/pii/S0957417424011813>

²⁰ https://www.mdpi.com/journal/entropy/special_issues/free_energy

https://www.mdpi.com/journal/entropy/special_issues/free_energy

²¹ <https://www.frontiersin.org/journals/network-physiology/articles/10.3389/fnetp.2025.1521963/full>

<https://www.frontiersin.org/journals/network-physiology/articles/10.3389/fnetp.2025.1521963/full>

²² <https://arxiv.org/abs/2505.19867>

<https://arxiv.org/abs/2505.19867>

²³ <https://arxiv.org/abs/2011.09756>

<https://arxiv.org/abs/2011.09756>

24 <https://www.mdpi.com/1099-4300/26/1/83>

<https://www.mdpi.com/1099-4300/26/1/83>

25 <https://arxiv.org/abs/2308.08307>

<https://arxiv.org/abs/2308.08307>

26 <https://arxiv.org/list/cs.MA/2022-10?show=50&skip=75>

<https://arxiv.org/list/cs.MA/2022-10?show=50&skip=75>

27 <https://www.sciencedirect.com/science/article/pii/S0149763423004694>

<https://www.sciencedirect.com/science/article/pii/S0149763423004694>

28 <https://arxiv.org/html/2411.07362v2>

<https://arxiv.org/html/2411.07362v2>

29 <https://arxiv.org/abs/2412.10425>

<https://arxiv.org/abs/2412.10425>

30 <https://www.frontiersin.org/journals/sustainable-cities/articles/10.3389/frsc.2025.1571613/full>

<https://www.frontiersin.org/journals/sustainable-cities/articles/10.3389/frsc.2025.1571613/full>

31 **Active Digital Twins via Active Inference**

https://arxiv.org/abs/2506.14453?utm_source=chatgpt.com