

# The Active Inference Notation Decoder

A Quick Reference Guide for Reading Across the Literature

Dr. Alianna J. Maren, Ph.D. | Themesis, Inc. | themesis.com

| The Rosetta Stone: Active Inference Notation Across Authors |                    |                  |                                                          |                                       |                       |
|-------------------------------------------------------------|--------------------|------------------|----------------------------------------------------------|---------------------------------------|-----------------------|
| Why the same concept looks different in every paper         |                    |                  |                                                          |                                       |                       |
| Concept                                                     | Namjoshi<br>(2026) | Beal<br>(2003)   | Friston<br>(Pre-2016)                                    | Friston<br>(2017 +)                   | Blei et al.<br>(2016) |
| Observed /<br>Observable                                    | $y$                | $y$              | $\tilde{s}, \tilde{a}$<br>(Markov blanket)               | $o$ or $y$                            | $x$                   |
| Hidden /<br>Latent                                          | $x$                | $x$              | $\tilde{\Psi}$<br>(External states)                      | $\eta$ (External<br>states)           | $z$                   |
| Internal /<br>Representational                              | -                  | -                | $\tilde{r}$                                              | -                                     | -                     |
| Approx. Posterior                                           | $q(x)$             | $q(x)$           | $q(\tilde{\Psi}   \tilde{r})$                            | $q(\eta   \mu)$                       | $q(z x)$              |
| Generative Model                                            | $p(y x)$           | $p(y x, \theta)$ | $p(\tilde{\Psi}, \tilde{s}, \tilde{a}, \tilde{r} m)$     | $p(o, \eta m)$                        | $p(x z)$              |
| "x" means ...                                               | hidden             | hidden           | " $\Psi$ " hidden<br>" $\tilde{s}, \tilde{a}$ " observed | " $\eta$ " hidden<br>" $o$ " observed | observed              |

## Why This Guide Exists

Active inference and variational inference are a single family of mathematical ideas — but reading across the literature feels like navigating without a map. The same concept appears under completely different symbols depending on which paper you are reading. A variable that means 'hidden state' in one paper means 'observed variable' in another. This is not a sign that you are confused. It is a sign that the field developed across multiple communities — statistical learning, theoretical neuroscience, and machine learning — each with its own notational conventions.

This guide is a Rosetta Stone: a single reference that maps the same concepts across five key sources. It accompanies the Themesis YouTube video *Surviving the Notational Nightmare: Active Inference and Generative AI Notation*.

## Key Definitions

**Observed / Observable ( $y, o, x, s, \tilde{a}$ ):** The variables that the agent can directly measure or sense. In neuroscience terms: sensory input. In machine learning terms: the data.

**Hidden / Latent ( $x, \Psi, \eta, z$ ):** The variables that the agent cannot directly observe but must infer. These are the hidden causes of observations. In Friston's framework, these are the external states of the world.

**Internal / Representational ( $r, \mu$ ):** Internal states of the agent that encode its beliefs. In Friston (pre-2016),  $r$  represents the internal representational states. In Friston (2017+),  $\mu$  (mu) represents the sufficient statistics of those

internal states — the compressed summary of the agent's beliefs.

**Approximate Posterior  $q(\cdot|\cdot)$ :** The agent's best tractable approximation to the true posterior over hidden states. Because the true posterior  $p(\text{hidden}|\text{observed})$  is computationally intractable, variational inference finds the closest approximation within a tractable family of distributions.

**Generative Model  $p(\cdot, \cdot|\cdot)$ :** The agent's model of how the world generates observations from hidden states. This encodes the agent's assumptions about the structure of reality. In pymdp, the A matrix encodes this likelihood mapping.

## The Critical Insight: What 'x' Means

In standard algebra and most of mathematics education,  $x$  is the observed or independent variable — the thing you measure, the input to a function.  $y$  depends on  $x$ .

In Namjoshi (2026) and Beal (2003), this is **reversed**:  $x$  is the **hidden** variable and  $y$  is what is **observed**. In Blei et al. (2016),  $x$  is observed and  $z$  is hidden — closer to the algebra convention but still different. In Friston's papers,  $x$  largely disappears as a primary variable after 2016.

This single collision between algebra intuition and active inference notation has derailed more serious readers than any other source of confusion in the literature.

## Primary References

- Namjoshi, S. (2026). *Fundamentals of Active Inference*. MIT Press. (January 2026)
- Beal, M.J. (2003). *Variational Algorithms for Approximate Bayesian Inference*. PhD Thesis, University College London.
- Friston, K.J. et al. (2017). Active Inference: A Process Theory. *Neural Computation*, 29(1), 1–49.
- Blei, D.M., Kucukelbir, A., & McAuliffe, J.D. (2017). Variational Inference: A Review for Statisticians. *Journal of the American Statistical Association*, 112(518), 859–877.
- Maren, A.J. (2019/2024). Derivation of the Variational Bayes Equations. Technical Report TR-2019-01, Themesis Inc. (Updated 2024)